



Cisco Nexus Dashboard Orchestrator Configuration Guide for ACI Fabrics, Release 4.0(x)

First Published: 2022-06-27

Last Modified: 2023-11-07

Americas Headquarters

Cisco Systems, Inc.
170 West Tasman Drive
San Jose, CA 95134-1706
USA
<http://www.cisco.com>
Tel: 408 526-4000
800 553-NETS (6387)
Fax: 408 527-0883

THE SPECIFICATIONS AND INFORMATION REGARDING THE PRODUCTS REFERENCED IN THIS DOCUMENTATION ARE SUBJECT TO CHANGE WITHOUT NOTICE. EXCEPT AS MAY OTHERWISE BE AGREED BY CISCO IN WRITING, ALL STATEMENTS, INFORMATION, AND RECOMMENDATIONS IN THIS DOCUMENTATION ARE PRESENTED WITHOUT WARRANTY OF ANY KIND, EXPRESS OR IMPLIED.

The Cisco End User License Agreement and any supplemental license terms govern your use of any Cisco software, including this product documentation, and are located at: <https://www.cisco.com/c/en/us/about/legal/cloud-and-software/software-terms.html>. Cisco product warranty information is available at <https://www.cisco.com/c/en/us/products/warranty-listing.html>. US Federal Communications Commission Notices are found here <https://www.cisco.com/c/en/us/products/us-fcc-notice.html>.

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 products and features described herein as in development or available at a future date remain in varying stages of development and will be offered on a when-and if-available basis. Any such product or feature roadmaps are subject to change at the sole discretion of Cisco and Cisco will have no liability for delay in the delivery or failure to deliver any products or feature roadmap items that may be set forth in this document.

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/c/en/us/about/legal/trademarks.html>. 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. (1721R)

© 2022 Cisco Systems, Inc. All rights reserved.



CONTENTS

CHAPTER 1	New and Changed Information 1
	New and Changed Information 1
PART I	Application & Fabric Management 3
CHAPTER 2	Tenants and Tenant Policies 5
	Tenants Overview 5
	Creating New Tenants 6
	Importing Existing Tenants 7
	Creating Tenant Policy Templates 8
CHAPTER 3	Schemas 17
	Schema and Template Design Considerations 17
	Shadow Objects 21
	Hiding Shadow Objects in APIC GUI 25
	Concurrent Configuration Updates 27
	Creating Schemas and Templates 29
	Importing Schema Elements From APIC Sites 30
	Configuring VRFs 31
	Configuring Bridge Domains 32
	Configuring Bridge Domain's Site-Local Properties 36
	Configuring Application Profiles and EPGs 37
	Configuring EPG's Site-Local Properties 39
	Configuring Contracts and Filters 42
	Configuring On-Premises External Connectivity 45
	Viewing Schemas 47

Bulk Update for Template Objects	47
Assigning Templates to Sites	50
Disassociating Template from Sites	50
Deploying Templates	51
Undeploying Templates	55
Template Versioning	56
Tagging Templates	56
Viewing History and Comparing Previous Versions	57
Reverting Template to Earlier Version	59
Template Review and Approval	60
Enabling Template Approval Requirement	60
Create Users with Required Roles	60
Requesting Template Review and Approval	61
Reviewing and Approving Templates	62
Configuration Drifts	63
Reconciling Configuration Drifts	64
Cloning Schemas	66
Cloning Templates	67
Migrating Objects Between Templates	68
Viewing Currently Deployed Configuration	69
Schema Overview and Deployment Visualizer	70

CHAPTER 4 **Fabric Management** 73

Creating Fabric Policies	73
Creating Fabric Resources Policies	84
Creating Monitoring Policies	88

PART II **Operations** 95

CHAPTER 5 **Audit Logs** 97

Audit Logs	97
------------	----

CHAPTER 6 **Backup and Restore** 99

Configuration Backup and Restore Guidelines	99
---	----

Configuring Remote Locations for Backups	101
Creating Backups	102
Restoring Backups	102
Exporting (Downloading) Backups	107
Importing Backups to Remote Location	108
Backup Scheduler	108

CHAPTER 7	Upgrading Sites	111
	Overview	111
	Guidelines and Limitations	113
	Downloading Controller and Switch Node Firmware to Sites	113
	Upgrading Controllers	116
	Upgrading Nodes	118

CHAPTER 8	Tech Support	123
	Tech Support and System Logs	123
	Downloading System Logs	124
	Streaming System Logs to External Analyzer	124

PART III	Infrastructure Management	129
-----------------	----------------------------------	------------

CHAPTER 9	System Configuration	131
	System Configuration Settings	131
	System Alias and Banner	131

CHAPTER 10	Preparing Cisco APIC Sites	133
	Pod Profile and Policy Group	133
	Configuring Fabric Access Policies for All APIC Sites	134
	Configuring Fabric Access Global Policies	134
	Configuring Fabric Access Interface Policies	135
	Configuring Sites That Contain Remote Leaf Switches	137
	Remote Leaf Guidelines and Limitations	137
	Configuring Routable Subnets for Remote Leaf Switches	137
	Enabling Direct Communication for Remote Leaf Switches	138

Cisco Mini ACI Fabrics 138

CHAPTER 11

Adding and Deleting Sites 141

Cisco NDO and APIC Interoperability Support 141

Adding Cisco ACI Sites 143

Removing Sites 144

Cross Launch to Fabric Controllers 145

CHAPTER 12

Configuring Infra General Settings 147

Infra Configuration Dashboard 147

Partial Mesh Intersite Connectivity 148

Configuring Infra: General Settings 149

CHAPTER 13

Configuring Infra for Cisco APIC Sites 153

Refreshing Site Connectivity Information 153

Configuring Infra: On-Premises Site Settings 154

Configuring Infra: Pod Settings 156

Configuring Infra: Spine Switches 157

CHAPTER 14

Configuring Infra for Cisco Cloud Network Controller Sites 161

Refreshing Cloud Site Connectivity Information 161

Configuring Infra: Cloud Site Settings 162

Recovering from Cloud Network Controller Site Downtime 164

CHAPTER 15

Deploying Infra Configuration for ACI Sites 167

Deploying Infra Configuration 167

Enabling Connectivity Between On-Premises and Cloud Sites 168

CHAPTER 16

CloudSec Encryption 173

Cisco ACI CloudSec Encryption 173

Requirements and Guidelines 174

CloudSec Encryption Terminology 177

CloudSec Encryption and Decryption Handling 178

CloudSec Encryption Key Allocation and Distribution	180
Configuring Cisco APIC for CloudSec Encryption	182
Configuring Cisco APIC for CloudSec Encryption Using GUI	183
Configuring Cisco APIC for CloudSec Encryption Using NX-OS Style CLI	183
Configuring Cisco APIC for CloudSec Encryption Using REST API	184
Enabling CloudSec Encryption in Nexus Dashboard Orchestrator	185
Verifying CloudSec Configuration on Switches	186
Rekey Process During Spine Switch Maintenance	188
Disabling and Re-Enabling Re-Key Process Using NX-OS Style CLI	188
Disabling and Re-Enabling Re-Key Process Using REST API	189

PART IV
Features and Use Cases 191

CHAPTER 17
DHCP Relay 193

DHCP Relay Policy	193
Guidelines and Limitations	193
Creating DHCP Relay Policies	194
Creating DHCP Option Policies	196
Assigning DHCP Policies	197
Creating DHCP Relay Contract	198
Verifying DHCP Relay Policies in APIC	199
Editing or Deleting Existing DHCP Policies	200

CHAPTER 18
EPG Preferred Group 201

EPG Preferred Groups Overview and Limitations	201
Configuring EPGs for Preferred Group	202

CHAPTER 19
Intersite L3Out 205

Intersite L3Out Overview	205
Intersite L3Out Guidelines and Limitations	206
Configuring External TEP Pool	207
Creating or Importing Intersite L3Out and VRF	208
Configuring External EPG to Use Intersite L3Out	210
Creating a Contract for Intersite L3Out	212

Use Cases	215
Intersite L3Out for Application EPGs (Intra-VRF)	215
Shared Services with Intersite L3Out for Application EPGs (Inter-VRF)	218
Intersite Transit Routing	221

CHAPTER 20

Intersite L3Out with PBR	225
Intersite L3Out with PBR	225
Supported Use Cases	226
Guidelines and Limitations	230
Configuring APIC Sites	230
Configuring External TEP Pool	230
Creating and Configuring L4-L7 Devices and PBR Policies	231
Creating Templates	234
Configuring Service Graph	236
Creating Filter and Contract	238
Creating Application EPG	244
Creating VRF and Bridge Domain for Application EPG	244
Creating Application Profile and EPG	245
Creating L3Out External EPG	247
Creating or Importing Intersite L3Out and VRF	247
Configuring External EPG	249

CHAPTER 21

Layer 3 Multicast	251
Layer 3 Multicast	251
Layer 3 Multicast Routing	252
Rendezvous Points	252
Multicast Filtering	253
Layer 3 Multicast Guidelines and Limitations	254
Creating Multicast Route Map Policy	255
Enabling Any-Source Multicast (ASM) Multicast	257
Enabling Source-Specific Multicast (SSM)	258

CHAPTER 22

QoS Preservation Across IPN	261
QoS and Global DSCP Policy	261

DSCP Policy Guidelines and Limitations	261
Configuring Global DSCP Policy	262
Set QoS Level for EPGs and Contracts	263

CHAPTER 23

SD-Access and ACI Integration 267

Cisco SD-Access and Cisco ACI Integration	267
Macro Segmentation	268
Cisco SD-Access and Cisco ACI Integration Guidelines	270
Onboarding the DNA Center	271
Configuring Connectivity Toward the SD-Access Domain	272
Viewing the Status of the SD-Access to ACI Integration	274
Extending a Virtual Network	276
Mapping or Unmapping a VN to a VRF	277
Configuring Transit Routing	279

CHAPTER 24

SD-WAN Integration 285

SD-WAN Integration	285
SD-WAN Integration Guidelines and Limitations	286
Adding a vManage Controller	287
Configuring Global DSCP Policy	288
Set QoS Level for EPGs and Contracts	290

CHAPTER 25

Multi-Site and SR-MPLS L3Out Handoff 293

Overview and Use Cases	293
SR-MPLS Infra Requirements and Guidelines	297
SR-MPLS Tenant Requirements and Guidelines	299
Creating Custom QoS Policy for SR-MPLS	301
Creating SR-MPLS Infra L3Out	303
Creating SR-MPLS Route Map Policy	306
Configure EPG-to-External-EPG (North-South) Communication	308

CHAPTER 26

vzAny Contracts 313

vzAny and Multi-Site	313
vzAny and Multi-Site Guidelines and Limitations	314

Create Contract and Filters	316
Configure vzAny to Consume/Provide a Contract	317
Create EPGs to Be Part of the vzAny VRF	317
Free Intra-VRF Communication	318
Stretched EPGs	320
Site-Local EPGs	320
Combination of Site-Local and Stretched EPGs	321
Intra-VRF Intersite L3Out	322
Inter-VRF Intersite L3Out	323
Many-to-One Communication	323
Provider EPG Within vzAny VRF	325
Provider EPG In Its Own VRF	325



CHAPTER 1

New and Changed Information

- [New and Changed Information](#), on page 1

New and Changed Information

The following table provides an overview of the significant changes to the organization and features in this guide from the release the guide was first published to the current release. The table does not provide an exhaustive list of all changes made to the guide.

Table 1: Latest Updates

Release	New Feature or Update	Where Documented
4.0(1)	First release of this document.	--



PART I

Application & Fabric Management

- [Tenants and Tenant Policies, on page 5](#)
- [Schemas, on page 17](#)
- [Fabric Management, on page 73](#)



CHAPTER 2

Tenants and Tenant Policies

- [Tenants Overview, on page 5](#)
- [Creating New Tenants, on page 6](#)
- [Importing Existing Tenants, on page 7](#)
- [Creating Tenant Policy Templates, on page 8](#)

Tenants Overview

A tenant is a logical container for application policies that enable an administrator to exercise domain-based access control. A tenant represents a unit of isolation from a policy perspective, but it does not represent a private network. Tenants can represent a customer in a service provider setting, an organization or domain in an enterprise setting, or just a convenient grouping of policies.

Three default tenants are pre-configured for you:

- `common`—A special tenant with the purpose of providing "common" services to other tenants in ACI fabrics. Global reuse is a core principle in the common tenant. Some examples of common services include shared L3Outs, DNS, DHCP, Active Directory, and shared private networks or bridge domains.
- `dcnm-default-tn`—A special tenant with the purpose of providing configuration for Cisco NDFC fabrics.
- `infra`—The Infrastructure tenant that is used for all internal fabric communications, such as tunnels and policy deployment. This includes switch to switch and switch to APIC communications. The `infra` tenant does not get exposed to the user space (tenants) and it has its own private network space and bridge domains. Fabric discovery, image management, and DHCP for fabric functions are all handled within this tenant.

When using Nexus Dashboard Orchestrator to manage Cisco NDFC fabrics, you will always use the default `dcnm-default-tn` tenant.



Note Nexus Dashboard Orchestrator cannot manage the APIC's `mgmt` tenant, so importing the tenant from APIC or creating a new tenant called `mgmt` in NDO it is not allowed.

To manage tenants, you must have either `Power User` or `Site and Tenant Manager` read-write role.

Tenant Policies Templates

Release 4.0(1) adds Tenant Policies templates, which allow you to configure the following tenant-wide policies:

- Route Policies for Multicast
- Route Map Policies for Route Control
- Custom QoS Policies
- DHCP Relay Policies
- DHCP Option Policies
- IGMP Interface Policies
- IGMP Snooping Policies
- MLD Snooping Policies

For additional information, see [Creating Tenant Policy Templates, on page 8](#).

Creating New Tenants

This section describes how to add a new tenant using the Nexus Dashboard Orchestrator GUI. If you want to import one or more existing tenants from your fabrics, follow the steps described in [Importing Existing Tenants, on page 7](#) instead.

Before you begin

You must have a user with either `Power User` or `Site Manager` read-write role to create and manage tenants.

Procedure

Step 1 Log in to your Nexus Dashboard and open the Nexus Dashboard Orchestrator service.

Step 2 Create a new tenant.

- a) From the left navigation pane, choose **Application Management > Tenants**.
- b) In the top right of the main pane, click **Add Tenant**.

The **Add Tenant** screen opens.

Step 3 Provide tenant details.

- a) Provide the **Display Name** and optional **Description**.

The tenant's **Display Name** is used throughout the Orchestrator's GUI whenever the tenant is shown. However, due to object naming requirements on the APIC, any invalid characters are removed and the resulting **Internal Name** is used when pushing the tenant to sites. The **Internal Name** that will be used when creating the tenant is displayed below the **Display Name** textbox.

Note

You can change the **Display Name** of the tenant at any time, but the **Internal Name** cannot be changed after the tenant is created.

- b) In the **Associated Sites** section, check all the sites you want to associate with this tenant.

Only the selected sites will be available for any templates using this tenant.

- c) (Optional) For each selected site, click the **Edit** button next to its name and choose one or more security domains.

A restricted security domain allows a fabric administrator to prevent a group of users, such as Tenant A, from viewing or modifying any objects created by a group of users in a different security domain, such as Tenant B, when users in both groups have the same assigned privileges. For example, a tenant administrator in Tenant A's restricted security domain will not be able to see policies, profiles, or users configured in Tenant B's security domain. Unless Tenant B's security domain is also restricted, Tenant B will be able to see policies, profiles, or users configured in Tenant A. Note that a user will always have read-only visibility to system-created configurations for which the user has proper privileges. A user in a restricted security domain can be given a broad level of privileges within that domain without the concern that the user could inadvertently affect another tenant's physical environment.

Security domains are created using the APIC GUI and can be assigned to various APIC policies and user accounts to control their access. For more information, see the *Cisco APIC Basic Configuration Guide*.

- d) In the **Associated Users** section, select the Nexus Dashboard Orchestrator users that are allowed to access the tenant.

Only the selected users will be able to use this tenant when creating templates.

Step 4 Click **Save** to finish adding the tenant.

Importing Existing Tenants

This section describes how to import one or more existing tenants. If you want to create a new tenant using Nexus Dashboard Orchestrator, follow the steps described in [Creating New Tenants, on page 6](#) instead.

Before you begin

You must have a user with either `Power User` or `Site Manager` read-write role to create and manage tenants.

Procedure

- Step 1** Log in to your Nexus Dashboard and open the Nexus Dashboard Orchestrator service.
- Step 2** In the left navigation menu, click **Sites**.
- Step 3** Locate the site from which you want to import the tenants, click its **Actions (...)** menu, and choose **Import Tenants**.
You can import tenants from one site at a time.
- Step 4** In the **Import Tenants** dialog, select one or more tenants to import and click **Ok**.
The selected tenants will be imported into the Nexus Dashboard Orchestrator and show in the **Application Management > Tenants** page.
- Step 5** Repeat these steps to import tenants from any other sites.
-

Creating Tenant Policy Templates

This section describes how to create one or more tenant policy templates. Tenant policy templates allow you to create and configure the following policies:

- Route Map Policies for Multicast
- Route Map Policies for Route Control
- Custom QoS Policies
- DHCP Relay Policies
- DHCP Option Policies
- IGMP Interface Policies
- MLD Snooping Policies

Procedure

Step 1 Log in to your Nexus Dashboard and open the Nexus Dashboard Orchestrator service.

Step 2 Create a new Tenant Policy template.

- From the left navigation pane, choose **Application Management > Tenant Policies**.
- On the **Tenant Policy Templates** page, click **Add Tenant Policy Template**.
- In the **Tenant Policies** page's right properties sidebar, provide the **Name** for the template.
- From the **Select a Tenant** dropdown, choose the tenant with which you want to associate this template.

All the policies you create in this template as described in the following steps will be associated with the selected tenant and deployed to it when you push the template to a specific site.

By default, the new template is empty, so you need to add one or more tenant policies as described in the following steps. Note that you don't have to create every policy available in the template – you can define one or more policies of each type to deploy along with this template. If you don't want to create a specific policy, simply skip the step that describes it.

Step 3 Assign the template to one or more sites.

The process for assigning Tenant Policy templates to sites is identical to how you assign application templates to sites.

- In the **Template Properties** view, click **Actions** and choose **Sites Association**.

The **Associate Sites to <template-name>** window opens.

- In the **Associate Sites** window, check the checkbox next to the sites where you want to deploy the template.

Note that only the on-premises ACI sites support tenant policy templates and will be available for assignment.

- Click **Ok** to save.

Step 4 Create a Route Map Policy for Multicast.

This policy is part of the overarching Layer 3 Multicast use case. You can use the information in this section as a reference, but we recommend following the full set of steps described in the [Layer 3 Multicast, on page 251](#) chapter of the *Features and Use Cases* section of this document.

- a) From the **+Create Object** dropdown, select **Route Map Policy for Multicast**.
- b) In the right properties sidebar, provide the **Name** for the policy.
- c) (Optional) Click **Add Description** and provide a description for the policy.
- d) Click **+Add Route Map for Multicast Entries** and provide the route map information.

For each route map, you need to create one or more route map entries. Each entry is a rule that defines an action based on one or more matching criteria based on the following information:

- **Order** – Order is used to determine the order in which the rules are evaluated.
- **Group IP, Src IP, and RP IP** – You can use the same multicast route map policy UI for two different use cases—to configure a set of filters for multicast traffic or to restrict a rendezvous point configuration to a specific set of multicast groups. Depending on which use case you're configuring, you only need to fill some of the fields in this screen:

- For multicast filtering, you can use the **Source IP** and the **Group IP** fields to define the filter. You must provide at least one of these fields, but can choose to include both. If one of the fields is left blank, it will match all values.

The Group IP range must be between 224.0.0.0 and 239.255.255.255 with a netmask between /4 and /32. You must provide the subnet mask.

The **RP IP** (Rendezvous Point IP) is not used for multicast filtering route maps, so leave this field blank.

- For Rendezvous Point configuration, you can use the **Group IP** field to define the multicast groups for the RP.

The Group IP range must be between 224.0.0.0 and 239.255.255.255 with a netmask between /4 and /32. You must provide the subnet mask.

For Rendezvous Point configuration, the **RP IP** is configured as part of the RP configuration. If a route-map is used for group filtering it is not necessary to configure an RP IP address in the route-map. In this case, leave the **RP IP** and **Source IP** fields empty.

- **Action** – Action defines the action to perform, either `Permit` or `Deny` the traffic, if a match is found.

- e) Click the checkmark icon to save the entry.
- f) Repeat the previous substeps to create any additional route map entries for the same policy.
- g) Click **Save** to save the policy and return to the template page.
- h) Repeat this step to create any additional Route Map for Multicast policies.

Step 5

Create a Route Map Policy for Route Control.

This policy is part of the overarching SR-MPLS use case. You can use the information in this section as a reference, but we recommend following the full set of steps described in the [Multi-Site and SR-MPLS L3Out Handoff, on page 293](#) chapter of the *Features and Use Cases* section of this document.

- a) From the **+Create Object** dropdown, select **Route Map Policy for Route Control**.
- b) In the right properties sidebar, provide the **Name** for the policy.
- c) (Optional) Click **Add Description** and provide a description for the policy.
- d) Click **+Add Entry** and provide the route map information.

For each route map, you need to create one or more context entries. Each entry is a rule that defines an action based on one or more matching criteria based on the following information:

- **Context Order** – Context order is used to determine the order in which contexts are evaluated. The value must be in the 0–9 range.
- **Context Action** – Context action defines the action to perform (permit or deny) if a match is found.

Once the context order and action are defined, choose how you want to match the context:

- Click **+Add Attribute** to specify the action that will be taken should the context match.

You can choose one of the following actions:

- Set Community
- Set Route Tag
- Set Dampening
- Set Weight
- Set Next Hop
- Set Preference
- Set Metric
- Set Metric Type
- Set AS Path
- Set Additional Community

After you have configured the attribute, click **Save**.

- If you want to match an action based on an IP address or prefix, click **Add IP Address**.

In the **Prefix** field, provide the IP address prefix. Both IPv4 and IPv6 prefixes are supported, for example 2003:1:1a5:1a5::/64 or 205.205.0.0/16.

If you want to aggregate IPs in a specific range, check the **Aggregate** checkbox and provide the range. For example, you can specify 0.0.0.0/0 prefix to match any IP or you can specify 10.0.0.0/8 prefix to match any 10.x.x.x addresses.

- If you want to match an action based on community lists, click **Add Community**.

In the **Community** field, provide the community string. For example, regular:as2-nn2:200:300.

Then choose the **Scope**: **Transitive** means the community will be propagated across eBGP peering (across autonomous systems) while **Non-Transitive** means the community will not be propagated.

- Repeat the previous substeps to create any additional route map entries for the same policy.
- Click **Save** to save the policy and return to the template page.
- Repeat this step to create any additional Route Map for Route Control policies.

Step 6

Create a Custom QoS Policy.

You can create a custom QoS policy in Cisco APIC to classify ingress traffic based on its DSCP or CoS values and associate it to a QoS priority level (QoS user class) to properly handle it inside the ACI fabric. Classification is supported

only if the DSCP values are present in the IP header and/or the CoS values are present in the Ethernet header of ingressing traffic. Additionally, the custom QoS policy can be used to modify the DSCP and/or CoS values in the header of ingressing traffic

As an example, custom QoS policies allow you to classify traffic coming into the ACI fabric traffic from devices that mark the traffic based only on the CoS value, such as Layer-2 packets which do not have an IP header.

For detailed information about QoS functionality in ACI fabrics, see [Cisco APIC and QoS](#).

- a) From the **+Create Object** dropdown, select **Custom QoS Policy**.
- b) In the right properties sidebar, provide the **Name** for the policy.
- c) (Optional) Click **Add Description** and provide a description for the policy.
- d) Click **+Add DSCP Mappings** and provide the required information.

The DSCP mapping configuration allows you to associate ingressing traffic, whose DSCP value is within the range specified in the mapping, to the specified QoS priority level (class). It also allows you to set the DSCP and/or CoS values of the ingressing traffic, so that those values can be retained when the traffic egresses the fabric.

Note

Retaining the target CoS value for egress traffic requires the configuration of the "Preserve CoS" policy, which is part of the NDO Fabric policies.

If the "DSCP Target" and/or "Target CoS" values are set as part of both the DSCP Mapping and CoS Mapping, the values specified in the DSCP Mapping have precedence.

For each mapping, you can specify the following fields:

- **DSCP From** – The start of the DSCP range.
- **DSCP To** – The end of the DSCP range.
- **DSCP Target** – The DSCP value to set on ingressing traffic that will be retained for egressing traffic.
- **Target CoS** – The CoS value to set on ingressing traffic that will be retained for egressing traffic when "Preserve CoS" is enabled.
- **Priority** – The QoS priority class to which the traffic will be assigned.

After you provide the mappings, click the checkmark icon to save. Then you can click **+Add DSCP Mappings** to provide additional mappings within the same policy.

- e) Click **Add** to save the policy and return to the template page.
- f) Click **+Add CoS Mappings** and provide the required information.

The DSCP mapping configuration allows you to associate ingressing traffic (whose DSCP value is within the range specified in the mapping) to the specified QoS priority level (class). It also allows you to set the DSCP and/or CoS values of the ingressing traffic, so that those values can be retained when the traffic egresses the fabric.

Note

Retaining the target CoS value for egress traffic requires the configuration of the "Preserve CoS" policy in the NDO Fabric policies.

In addition, if the "DSCP Target" and/or "Target CoS" values are set as part of both the DSCP Mapping and CoS Mapping, the values specified in the DSCP Mapping have precedence.

For each mapping, you can specify the following fields:

- **Dot1P From** – The start of the CoS range.

- **Dot1P To** – The end of the CoS range.
- **DSCP Target** – The DSCP value to set on ingress traffic that will be retained for egress traffic.
- **Target CoS** – The CoS value to set on ingress traffic that will be retained for egress traffic when "Preserve CoS" is enabled.
- **Priority** – The QoS priority class to which the traffic will be assigned.

After you provide the mappings, click the checkmark icon to save. Then you can click **+Add Cos Mappings** to provide additional mappings within the same policy.

- Click **Add** to save the policy and return to the template page.
- Repeat this step to create any additional Route Map for Route Control policies.

Step 7

Create a DHCP Relay Policy.

This policy is part of the overarching DHCP Relay use case. You can use the information in this section as a reference, but we recommend following the full set of steps described in the [DHCP Relay, on page 193](#) chapter of the *Features and Use Cases* section of this document.

- From the **+Create Object** dropdown, select **DHCP Relay Policy**.
- In the right properties sidebar, provide the **Name** for the policy.
- (Optional) Click **Add Description** and provide a description for the policy.
- Click **Add Provider** to configure the DHCP server to which you want to relay the DHCP requests originated by the endpoints.
- Select the provider type.

When adding a relay policy, you can choose one of the following two types:

- **Application EPG**—specifies the application EPG that includes the DHCP server to which you want to relay the DHCP requests.
- **L3 External Network**—specifies the External EPG associated to the L3Out that is used to access the network external to the fabric where the DHCP server is connected.

Note

You can select any EPG or external EPG that has been created in the Orchestrator and assigned to the tenant you specified, even if you have not yet deployed it to sites. If you select an EPG that hasn't been deployed, you can still complete the DHCP relay configuration, but you will need to deploy the EPG before the relay is available for use.

- Click **Select an Application EPG** or **Select an External EPG** (based on the provider type you selected) and choose the provider EPG.
- In the **DHCP Server Address** field, provide the IP address of the DHCP server.
- Enable the **DHCP Server VRF Preference** option if required.

This feature was introduced in Cisco APIC release 5.2(4). For more information on the use cases where it is required refer to the [Cisco APIC Basic Configuration Guide](#).

- Click **OK** to save the provider information.
- Repeat the previous substeps for any additional providers in the same DHCP Relay policy.
- Repeat this step to create any additional DHCP Relay policies.

Step 8

Create a DHCP Option Policy.

This policy is part of the overarching DHCP Relay use case. You can use the information in this section as a reference, but we recommend following the full set of steps described in the [DHCP Relay, on page 193](#) chapter of the *Features and Use Cases* section of this document.

- a) From the **+Create Object** dropdown, select **DHCP Option Policy**.
- b) In the right properties sidebar, provide the **Name** for the policy.
- c) (Optional) Click **Add Description** and provide a description for the policy.
- d) Click **Add Option**.
- e) Provide option details.

For each DHCP option, provide the following:

- **Name** – While not technically required, we recommend using the same name for the option as listed in [RFC 2132](#).

For example, `Name Server`.

- **Id** – Provide the value if the option requires one.

For example, a list of name servers available to the client for the Name Server option.

- **Data** – Provide the value if the option requires one.

For example, a list of name servers available to the client for the Name Server option.

- f) Click **OK** to save.
- g) Repeat the previous substeps for any additional options in the same DHCP Option policy.
- h) Repeat this step to create any additional DHCP Option policies.

Step 9

Create a IGMP Interface Policy.

IGMP snooping examines IP multicast traffic within a bridge domain to discover the ports where interested receivers reside. Using the port information, IGMP snooping can reduce bandwidth consumption in a multi-access bridge domain environment to avoid flooding the entire bridge domain.

For detailed information on IGMP snooping in ACI fabrics, see the "IGMP Snooping" chapter of the [Cisco APIC Layer 3 Networking Configuration Guide](#) for your release.

- a) From the **+Create Object** dropdown, select **IGMP Interface Policy**.
- b) In the right properties sidebar, provide the **Name** for the policy.
- c) (Optional) Click **Add Description** and provide a description for the policy.
- d) Provide policy details.

- **Allow Version 3 ASM** – Allow accepting IGMP version 3 source-specific reports for multicast groups outside of the SSM range. When this feature is enabled, the switch will create an (S,G) mroute entry if it receives an IGMP version 3 report that includes both the group and source even if the group is outside of the configured SSM range. This feature is not required if hosts send (*,G) reports outside of the SSM range, or send (S,G) reports for the SSM range

- **Fast Leave** – Option that minimizes the leave latency of IGMPv2 group memberships on a given IGMP interface because the device does not send group-specific queries. When Fast Leave is enabled, the device removes the group entry from the multicast routing table immediately upon receiving a leave message for the group. The default is disabled.

Use this only when there is only one receiver behind the BD/interface for a given group.

- **Report Link Local Groups** – Enables sending reports for groups in 224.0.0.0/24. Reports are always sent for nonlink local groups. By default, reports are not sent for link local groups.

- **IGMP Version** – IGMP version that is enabled on the bridge domain or interface. The IGMP version can be 2 or 3. The default is 2.
- **Advanced Settings** – Click the arrow next to this section to expand.
 - **Group Timeout** – Group membership interval that must pass before the router decides that no members of a group or source exist on the network. Values range from 3 to 65,535 seconds. The default is 260 seconds.
 - **Query Interval** – Sets the frequency at which the software sends IGMP host query messages. Values can range from 1 to 18,000 seconds. The default is 125 seconds.
 - **Query Response Interval** – Sets the response time advertised in IGMP queries. Values can range from 1 to 25 seconds. The default is 10 seconds.
 - **Last Member Count** – Sets the number of times that the software sends an IGMP query in response to a host leave message. Values can range from 1 to 5. The default is 2.
 - **Last Member Response Time** – Sets the query interval waited after sending membership reports before the software deletes the group state. Values can range from 1 to 25 seconds. The default is 1 second.
 - **Startup Query Count** – Configures snooping for a number of queries sent at startup when you do not enable PIM because multicast traffic does not need to be routed. Values can range from 1 to 10. The default is 2 messages.
 - **Startup Query Interval** – This configures the IGMP snooping query interval at startup. The range is from 1 second to 18000 seconds. The default is 125 seconds..
 - **Querier Timeout** – Sets the query timeout that the software uses when deciding to take over as the querier. Values can range from 1 to 65,535 seconds. The default is 255 seconds
 - **Robustness Variable** – Sets the robustness variable. You can use a larger value for a lossy network. Values can range from 1 to 7. The default is 2.
 - **State Limit Route Map** – Used with Reserved Multicast Entries feature.
The route map policy must be already created as described in Step 2.
 - **Report Policy Route Map** – Access policy for IGMP reports that is based on a route-map policy. IGMP group reports will only be selected for groups allowed by the route-map.
The route map policy must be already created as described in Step 2.
 - **Static Report Route Map** – Statically binds a multicast group to the outgoing interface, which is handled by the switch hardware. If you specify only the group address, the (*, G) state is created. If you specify the source address, the (S, G) state is created. You can specify a route-map policy name that lists the group prefixes, group ranges, and source prefixes. A source tree is built for the (S, G) state only if you enable IGMPv3.
The route map policy must be already created as described in Step 2.
 - **Maximum Multicast Entries** – Limit the mroute states for the BD or interface that are created by IGMP reports. Default is disabled and no limit is enforced. Valid range is 1-4294967295.

e) Repeat this step to create any additional IGMP Interface policies.

Step 10

Create a MLD Snooping Policy.

Multicast Listener Discovery (MLD) snooping enables the efficient distribution of IPv6 multicast traffic between hosts and routers. It is a Layer 2 feature that restricts IPv6 multicast traffic within a bridge domain to a subset of ports that have transmitted or received MLD queries or reports. In this way, MLD snooping provides the benefit of conserving the bandwidth on those segments of the network where no node has expressed interest in receiving the multicast traffic. This reduces the bandwidth usage instead of flooding the bridge domain, and also helps hosts and routers save unwanted packet processing.

For detailed information on MLD snooping in ACI fabrics, see the "MLD Snooping" chapter of the [Cisco APIC Layer 3 Networking Configuration Guide](#) for your release.

- a) From the **+Create Object** dropdown, select **MLD Snooping Policy**.
- b) In the right properties sidebar, provide the **Name** for the policy.
- c) (Optional) Click **Add Description** and provide a description for the policy.
- d) Provide policy details.

- **Admin State** – Enables or disables the MLD snooping feature.

- **Fast Leave Control** – Allows you to turn on or off the fast-leave feature on a per bridge domain basis. This applies to MLDv2 hosts and is used on ports that are known to have only one host doing MLD behind that port.

Default is `disabled`.

- **Querier Control** – Enables or disables MLD snooping querier processing. MLD snooping querier supports the MLD snooping in a bridge domain where PIM and MLD are not configured because the multicast traffic does not need to be routed

Default is `disabled`.

- **Querier Version** – Allows you to choose the querier version.

Default is `version2`.

- **Advanced Settings** – Click the arrow next to this section to expand.

- **Query Interval** – Sets the frequency at which the software sends MLD host query messages. Values can range from 1 to 18,000 seconds.

The default is 125 seconds.

- **Query Response Interval** – Sets the response time advertised in MLD queries. Values can range from 1 to 25 seconds.

The default is 10 seconds.

- **Last Member Query Interval** – Sets the query response time after sending membership reports before the software deletes the group state. Values can range from 1 to 25 seconds.

The default is 1 second.

- **Start Query Count** – Configures snooping for a number of queries sent at startup when you do not enable PIM because multicast traffic does not need to be routed. Values can range from 1 to 10.

The default is 2.

- **Start Query Interval** – Configures a snooping query interval at startup when you do not enable PIM because multicast traffic does not need to be routed. Values can range from 1 to 18,000 seconds.

The default is 31 seconds.

- e) Repeat this step to create any additional MLD Snooping policies.

Step 11 Click **Save** to save the changes you've made to the template.

Note

When you save (or deploy) the template to one or more sites, the Orchestrator will verify that the specified nodes and/or interfaces are valid for the site(s) and will return an error.

Step 12 Click **Deploy** to deploy the template to the associated sites.

The process for deploying tenant policy templates is identical to how you deploy application templates.

If you have previously deployed this template but made no changes to it since, the **Deploy** summary will indicate that there are no changes and you can choose to re-deploy the entire template. In this case, you can skip this step.

Otherwise, the **Deploy to sites** window will show you a summary of the configuration differences that will be deployed to sites. Note that in this case only the difference in configuration is deployed to the sites. If you want to re-deploy the entire template, you must deploy once to sync the differences and then redeploy again to push the entire configuration as described in the previous paragraph.



CHAPTER 3

Schemas

- [Schema and Template Design Considerations, on page 17](#)
- [Shadow Objects, on page 21](#)
- [Concurrent Configuration Updates, on page 27](#)
- [Creating Schemas and Templates, on page 29](#)
- [Bulk Update for Template Objects, on page 47](#)
- [Assigning Templates to Sites, on page 50](#)
- [Deploying Templates, on page 51](#)
- [Template Versioning, on page 56](#)
- [Template Review and Approval, on page 60](#)
- [Configuration Drifts, on page 63](#)
- [Cloning Schemas, on page 66](#)
- [Cloning Templates, on page 67](#)
- [Migrating Objects Between Templates, on page 68](#)
- [Viewing Currently Deployed Configuration, on page 69](#)
- [Schema Overview and Deployment Visualizer, on page 70](#)

Schema and Template Design Considerations

Nexus Dashboard Orchestrator provides a number of policy templates that allow you to define one or more policies together and deploy them to one or more sites at the same time. These include tenant policy templates, fabric and fabric resources policy templates, monitoring templates, and application templates. A schema is a collection of application templates, which are used for defining application policies, with each template assigned to a specific tenant. There are multiple approaches you can take when it comes to creating schema and application template configurations specific to your deployment use case. The following sections describe a few simple design directions you can take when deciding how to define the schemas, templates, and policies in your Multi-Site environment.

Keep in mind that when designing schemas, you must consider the supported scalability limits for the number of schemas, templates, and objects per schema. Detailed information on verified scalability limits is available in the [Cisco Multi-Site Verified Scalability Guides](#) for your release.

Application Templates

There are 3 types of schema templates, also known as application templates, available in Nexus Dashboard Orchestrator, each designed for a specific purpose:

- **ACI Multi-Cloud**—Templates used for Cisco ACI on-premises and cloud sites. This template supports two deployment types:

- **Multi-Site** - The template can be associated to a single site (site-local policies) or to multiple sites (stretched policies) and the option should be selected for Multi-Site Network (ISN) or VXLAN intersite communication to allow template and object stretching between multiple sites.
- **Autonomous** - The template can be associated to one or more sites that are operated independently and are not connected through an Inter-Site Network (no intersite VXLAN communication).

Because autonomous sites are by definition isolated and do not have any intersite connectivity, there is no shadow object configuration across sites and no cross-programming of pctxs or VNIDs in the spine switches for intersite traffic flow.

The autonomous templates also allow for significantly higher deployment scale.

The following sections focus primarily on this type of templates.

- **NDFC**—Templates designed for Cisco Nexus Dashboard Fabric Controller (formerly Data Center Network Manager) sites.

This guide described Nexus Dashboard Orchestrator configurations for on-premises Cisco ACI fabrics. For information on working with Cisco NDFC sites, see the [Cisco Nexus Dashboard Orchestrator Configuration Guide for NDFC Fabrics](#) instead.

- **Cloud Local**—Templates designed for specific Cloud Network Controller use cases, such as Google Cloud site connectivity, and cannot be stretched between multiple sites.

This guide described Nexus Dashboard Orchestrator configurations for on-premises Cisco ACI fabrics. For information on working with Cloud Network Controller fabrics, see the [Nexus Dashboard Orchestrator use case library](#) instead.

When creating schemas and application templates, you can choose to adopt one of the following simple approaches:

- **Single Template Deployment**

The simplest schema design approach is a single schema, single template deployment. You can create a single schema with a single template within it and add all VRFs, Bridge Domains, EPGs, Contracts and other elements to that template and deploy it to one or more sites.

This simplest approach to Multi-Site schema creation is to create all objects within the same schema and template. However, the supported number of schemas scalability limit may make this approach unsuitable for large scale deployments, which could exceed those limits.

Note also that with this approach all the objects defined in the template become "stretched objects" and all changes made to the template are always simultaneously deployed to all the sites associated to such template.

- **Multiple Templates with Network Separation**

Another approach to schema design is to separate the networking objects from the application policy configuration. Networking objects include VRFs, Bridge Domains, and subnets, while the application policy objects include EPGs, Contracts, Filters, External EPGs, and Service Graphs.

You begin by defining a schema that contains the network elements. You can choose to create a single schema that contains all the network elements or you can split them into multiple schemas based on which applications reference them or which sites the network is stretched to.

You can then define one or more separate schemas which contain each application's policy objects. This new schema can reference the network elements, such as bridge domains, defined in the previous schema.

After creating and deploying the policy schemas and templates, the networking objects in the networking schema will display the number of external references by the policy schema elements. The object with external references will also be denoted by the ribbon icon.

Schemas designed this way provide logical separation of networking objects from the policy objects. However, this creates additional complexity when it comes to keeping track of externally referenced objects in each schema.

• Multiple Templates Based On Object Relationships

When configuring multiple schemas with shared object references, it is important to be careful when making changes to those objects. For instance, making changes to or deleting a shared networking object can impact applications in one or more sites. Because of that, you may choose to create a template around each individual site that contains only the objects used by that site and its applications, including the VRFs, BDs, EPGs, Contracts, and Filters. And create different templates containing the shared objects.

For example, you can create a `Site1` template that contains only the objects that are local to Site1 and the template is deployed to only that site. Similarly, the `Site2` template contains only the object relevant to site2 and is deployed to that site only. Any change made to any object in either of these templates has no effect on the other one. Then you can create a `shared` template which contains objects that are shared between the sites.

You can extend this scenario for an additional site with the following template layout:

- Site 1 template
- Site 2 template
- Site 3 template
- Site 1 and 2 shared template
- Site 1 and 3 shared template
- Site 2 and 3 shared template
- All shared template

Similarly, rather than separating objects based on which site they are deployed to, you can also choose to create schemas and templates based on individual applications instead. This would allow you to easily identify each application profile and map them to schemas and sites as well as easily configure each application as local or stretched across sites.

However, as this could quickly exceed the templates per schema limit (listed in the [Verified Scalability Guide](#) for your release), you would have to create additional schemas to accommodate the multiple combinations. While this creates additional complexity with multiple additional schemas and templates, it provides true separation of objects based on site or application.

Fabric Policy Templates

In addition to the three types of application templates, Release 4.0(1) adds 3 new templates designed for fabric-wide policies:

- **Fabric Policies** templates can be used for managing the following fabric-wide policies:

- VLAN Pool
- Physical Domains
- SyncE Interface Policies
- Interface Settings
- Node Settings
- Pod Settings
- MACsec
- NTP Policies
- PTP Policies
- QoS DSCP Policies
- QoS SR-MPLS Policies
- QoS Class Policies

For additional information, see [Creating Fabric Policies, on page 73](#).

- **Fabric Resources Policies** templates can be used for managing the following fabric-wide policies:

- Physical Interfaces
- Port Channel Interfaces
- Virtual Port Channel Interfaces
- Node Profiles

These templates reference policies defined in the Fabric Policies templates, so those templates must be created and deployed first. For additional information, see [Creating Fabric Resources Policies, on page 84](#).

- **Monitoring Policy** templates can be used for managing `Tenant SPAN` or `Access SPAN` policies.

For additional information, see [Creating Monitoring Policies, on page 88](#).

Template Design Best Practices

Beginning with Release 4.0(1), Nexus Dashboard Orchestrator will validate and enforce a number of best practices when it comes to template design and deployment. Regardless of the type of template you are creating, keep in mind the following:

- All policy objects must be **deployed** in order according to their dependencies.

For example, when creating a bridge domain (BD), you must associate it with a VRF. In this case, the BD has a VRF dependency so the VRF must be deployed to the fabric before or together with the BD. If these two objects are defined in the same template, then the Orchestrator will ensure that during deployment, the VRF is created first and associate it with the bridge domain.

However, if you define these two objects in separate templates and attempt to deploy the template with the BD first, the Orchestrator will return a validation error as the associated VRF is not yet deployed. In this case you must deploy the VRF template first, followed by the BD template.

- All policy objects must be **undeployed** in order according to their dependencies, or in the opposite order in which they were deployed.

As a corollary to the point above, when you undeploy templates, you must not undeploy objects on which other objects depend. For example, you cannot undeploy a VRF before undeploying the BD with which the VRF is associated.

- No cyclical dependencies are allowed across multiple templates.

Consider a case of a VRF (`vrf1`) associated with a bridge domain (`bd1`), which is in turn associated with an EPG (`epg1`). If you create `vrf1` in `template1` and deploy that template, then create `bd1` in `template2` and deploy that template, there will be no validation errors since the objects are deployed in correct order. However, if you then attempt to create `epg1` in `template1`, it would create a circular dependency between the two template, so the Orchestrator will not allow you to save `template1` addition of the EPG.

Shadow Objects

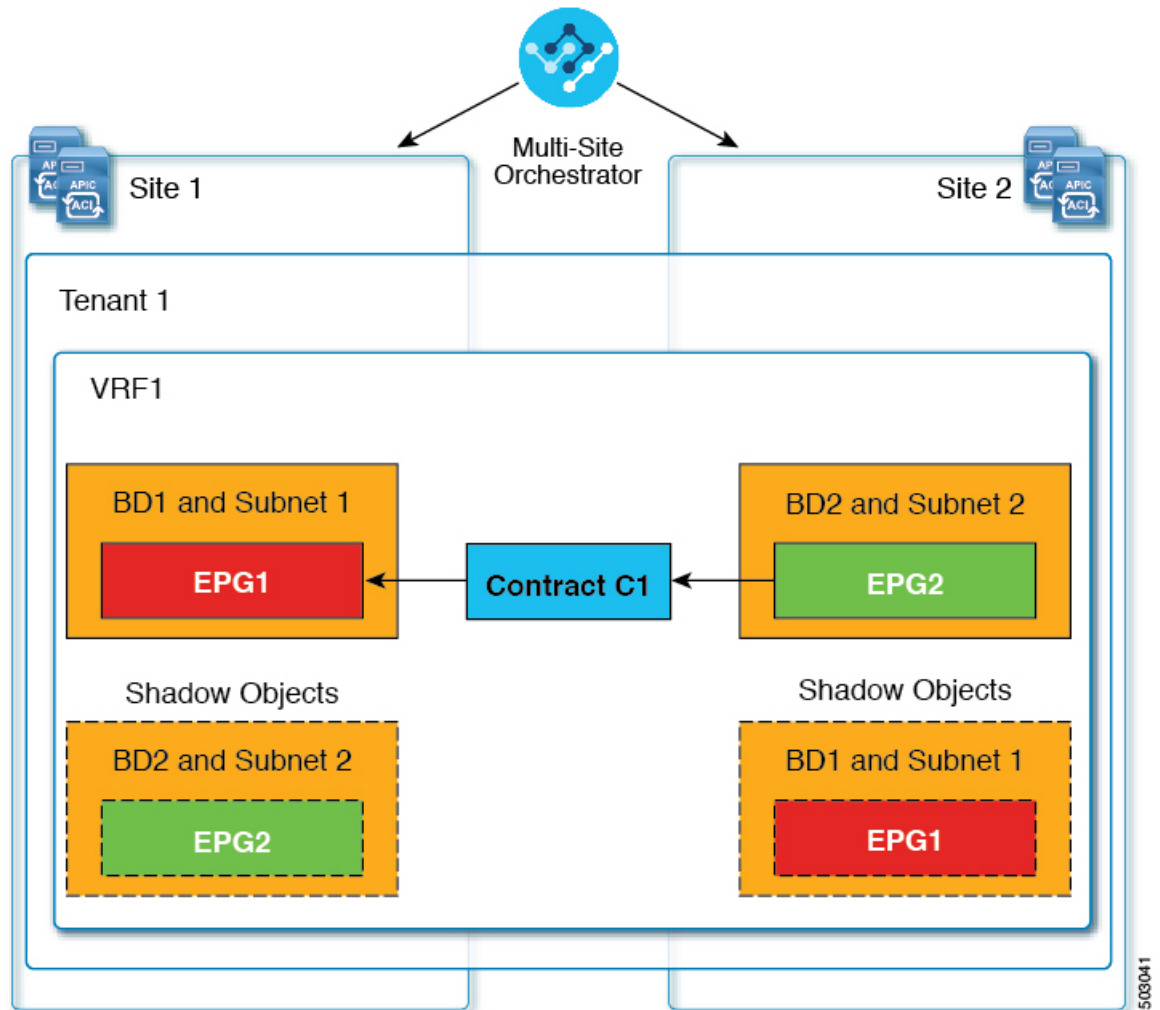
When a contract exists between site-local EPGs in stretched VRF or in Shared Services use-cases where provider and consumer are in different VRFs and communicate through Tenant contracts, the EPGs and bridge domains (BDs) are mirrored on the remote sites. The mirrored objects appear as if they are deployed in each of these sites' controllers, while only actually being deployed in one of the sites. These mirrored objects are called "shadow" objects.



Note Shadow objects should not be removed using the APIC GUI.

For example, if a tenant and VRF are stretched between Site1 and Site2, provider EPG and its bridge domain are deployed in Site2 only, and consumer EPG and its domain are deployed in Site1 only, then corresponding shadow bridge domains and EPGs will be deployed as shown in the figure below. They appear with the same names as the ones that were deployed directly to each site.

Figure 1: Basic Shadow EPG



The following objects can be shadowed:

- VRFs
- Bridge Domains (BDs)
- L3Outs
- External EPGs
- Application Profiles
- Application EPGs
- Contracts (Hybrid Cloud deployments)

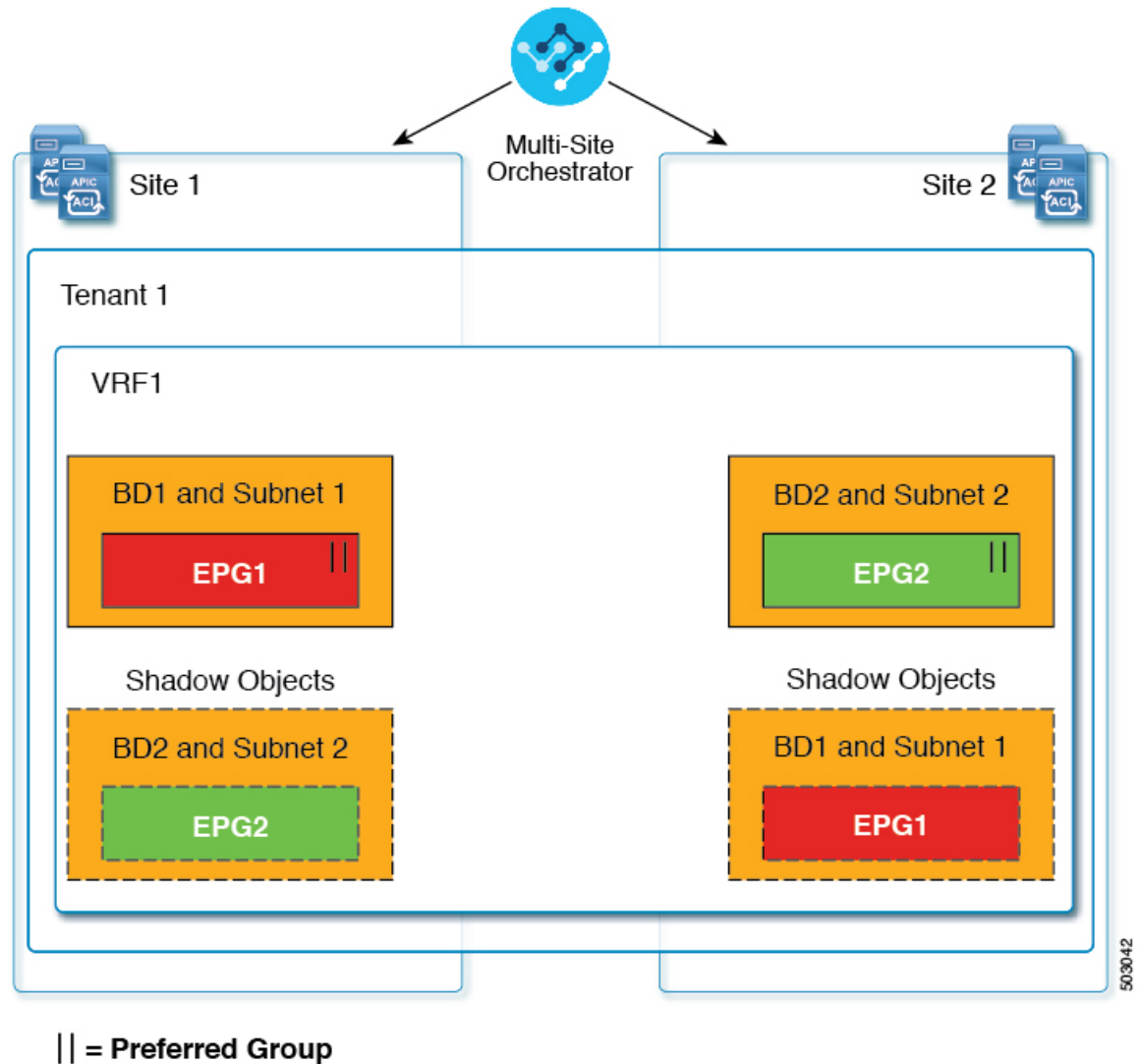
If your fabrics are running APIC Release 5.0(2) or later, when you select a shadow object in the APIC GUI, you will see a **This is a shadow object pushed by MSC to support intersite policies. Do not make any changes or delete this object.** warning at the top of main GUI pane. In addition, shadow

EPGs that are not part of a VMM domain will not have static ports, while shadow BDs will have **No Default SVI Gateway** option enabled in the APIC GUI.

Other Use Cases with Shadow Objects

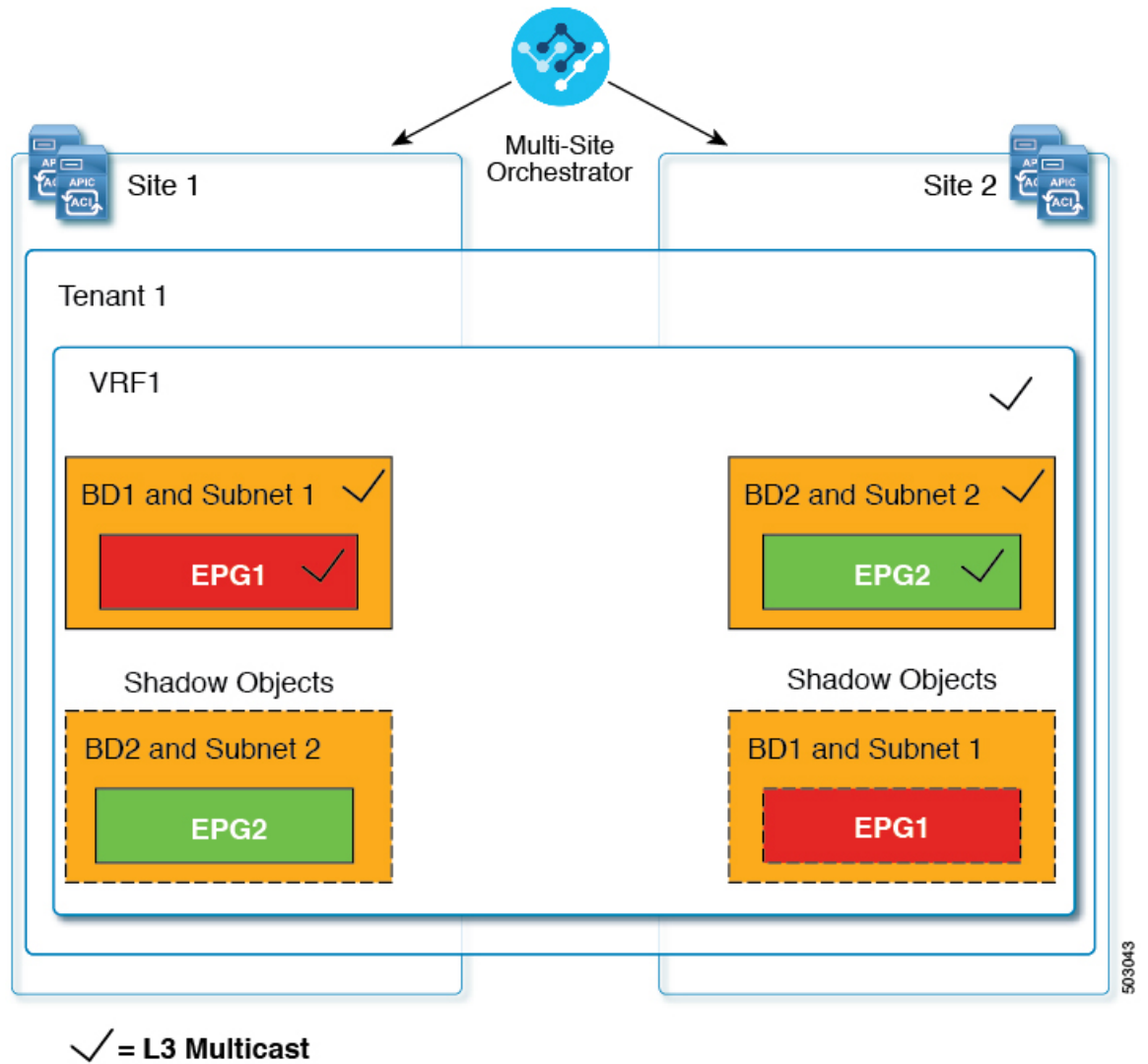
Shadow objects are also create in a number of other use cases, such as Preferred Group, vzAny, and Layer 3 Multicast, and hybrid cloud, as shown in the figures below.

Figure 2: Preferred Group



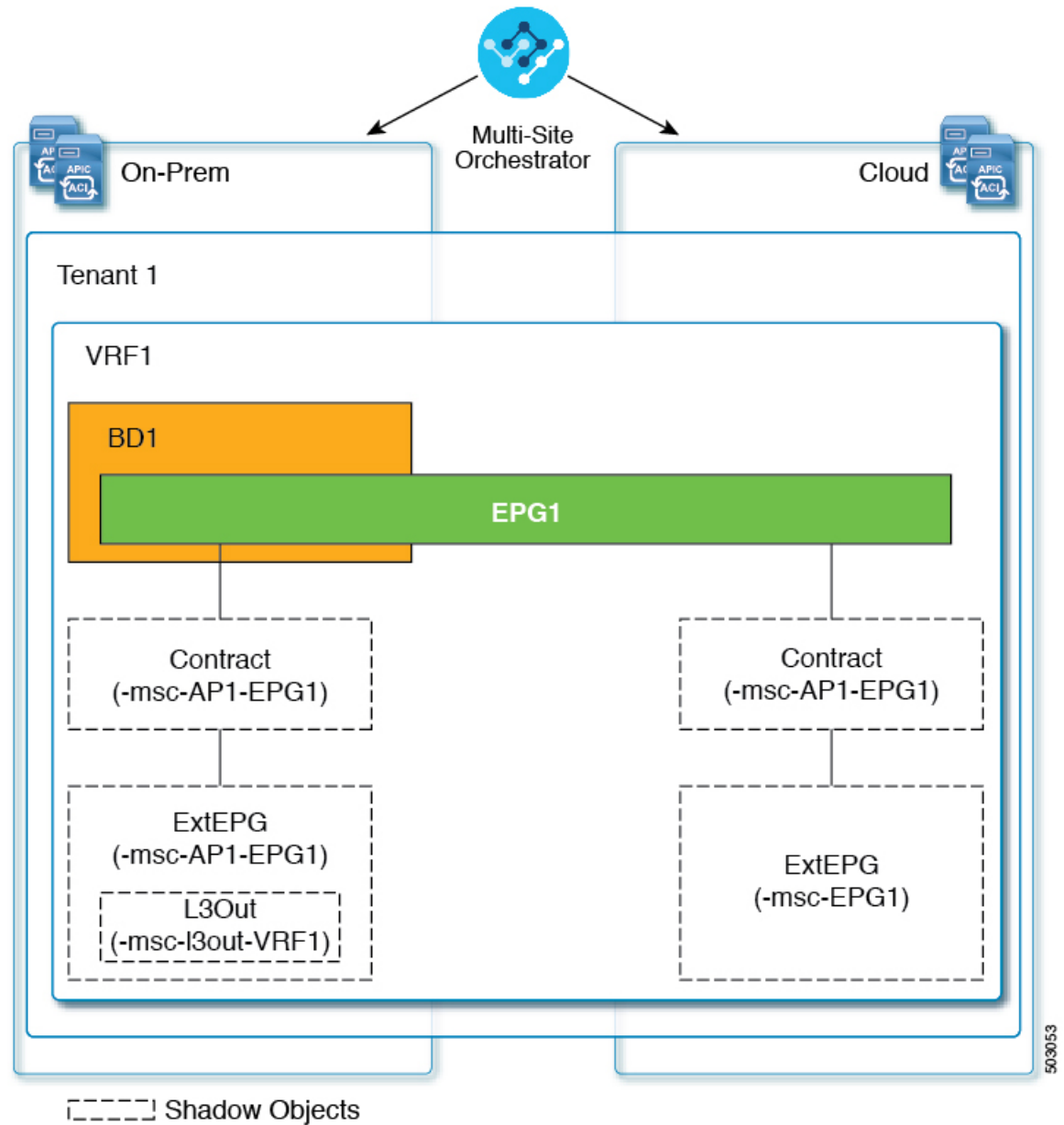
In case of multicast, the shadow objects are created only for EPGs/BDs that have multicast sources connected and the option explicitly configured at the EPG level.

Figure 3: L3 Multicast



In case of hybrid cloud deployments, even stretched objects will create shadow objects where implicit contracts exist. For example in the following case where an EPG is stretched between an on-premises and cloud sites, shadow external EPGs are created in each site with implicit shadow contracts between the stretched EPG and the shadow external EPGs.

Figure 4: Hybrid Cloud



Starting with Cisco APIC, Release 5.2(3), shadow objects are indicated by a unique icon in the Cisco APIC GUI. Regular Orchestrator-created objects are shown with a green cloud symbol, whereas the shadow objects will have a gray cloud icon.

Hiding Shadow Objects in APIC GUI

Starting with APIC Release 5.0(2), you can choose to show or hide the shadow objects created by the Nexus Dashboard Orchestrator in the on-premises site's APIC GUI. Shadow objects in Cloud Network Controller are always hidden.

If you want to hide shadow objects from the GUI, keep the following in mind:

- This option cannot be set globally from the Orchestrator and must be set directly in each site's APIC as described in this section.
- The option to show shadow objects is turned off by default for all new APIC Release 5.0(2) installations and upgrades, so previously visible objects may become hidden.
- Hiding shadow objects relies on a flag set by the Nexus Dashboard Orchestrator specifically for this feature, which is enabled from Orchestrator Release 3.0(2) and later:
 - If shadow objects are deployed by an earlier Orchestrator version, they will not have the required tag and will always be visible in the APIC GUI.
 - If shadow objects are deployed by Orchestrator version 3.0(2) or later, they will have the tag and can be hidden or shown using the APIC GUI setting.
 - We recommend upgrading each fabric to APIC Release 5.0(2) before upgrading the Nexus Dashboard Orchestrator.

When the Nexus Dashboard Orchestrator is upgraded to Release 3.0(2), any objects deployed to sites running APIC Release 5.0(2) or later will be tagged with appropriate tags and can be shown or hidden using the APIC GUI without having to re-deploy them.

If you upgrade the Orchestrator before the fabric's APIC, the site's objects will not be tagged and you will need to manually re-deploy the configuration after the fabric is upgraded for the flag to be set.

- If you ever downgrade your fabric to a release prior to Release 5.0(2), the shadow objects will no longer be hidden and you may see a different icon for them in the APIC GUI.



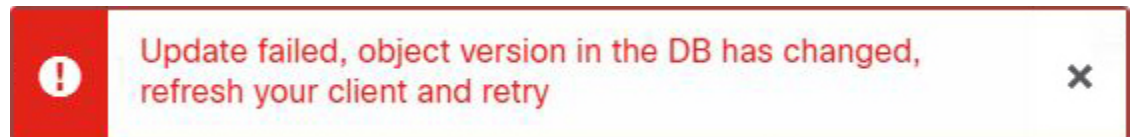
Procedure

-
- Step 1** Log in to the site's APIC.
- Step 2** In the top right corner, click the **Manage my profile** icon and choose **Settings**.
- Step 3** In the **Application Settings** window, enable or disable the **Show Hidden Policies** checkbox.
- The setting is stored in the user profile and is enable or disabled separately for each user.

Step 4 Repeat the process for any additional APIC sites.

Concurrent Configuration Updates

The Nexus Dashboard Orchestrator GUI will ensure that any concurrent updates on the same site or schema object cannot unintentionally overwrite each other. If you attempt to make changes to a site or template that was updated by another user since you opened it, the GUI will reject any subsequent changes you try to make and present a warning requesting you to refresh the object before making additional changes; refreshing the template will lose any edits you made up to that point and you will have to make those changes again:



However, the default REST API functionality was left unchanged in order to preserve backward compatibility with existing applications. In other words, while the UI is always enabled for this protection, you must explicitly enable it for your API calls for NDO to keep track of configuration changes.



Note When enabling this feature, note the following:

- This release supports detection of conflicting configuration changes for Site and Schema objects only.
- Only `PUT` and `PATCH` API calls support the version check feature.
- If you do not explicitly enable the version check parameter in your API calls, NDO will not track any updates internally. And as a result, any configuration updates can be potentially overwritten by both subsequent API calls or GUI users.

To enable the configuration version check, you can pass the `enableVersionCheck=true` parameter to the API call by appending it to the end of the API endpoint you are using, for example:

```
https://<mso-ip-address>/mso/api/v1/schemas/<schema-id>?enableVersionCheck=true
```

Example

We will use a simple example of updating the display name of a template in a schema to show how to use the version check attribute with `PUT` or `PATCH` calls.

First, you would `GET` the schema you want to modify, which will return the current latest version of the schema in the call's response:

```
{
  "id": "601acfed38000070a4ee9ec0",
  "displayName": "Schema1",
  "description": "",
  "templates": [
    {
      "name": "Template1",
      "displayName": "current name",
      [...]
    }
  ]
}
```

```

    }
  ],
  "_updateVersion": 12,
  "sites": [...]
}

```

Then you can modify the schema in one of two ways appending `enableVersionCheck=true` to the request URL:



Note You must ensure that the value of the `"_updateVersion"` field in the payload is the same as the value you got in the original schema.

- Using the `PUT` API with the entire updated schema as payload:

```

PUT /v1/schemas/601acfed38000070a4ee9ec0?enableVersionCheck=true

{
  "id": "601acfed38000070a4ee9ec0",
  "displayName": "Schema1",
  "description": "",
  "templates": [
    {
      "name": "Template1",
      "displayName": "new_name",
      [...]
    }
  ],
  "_updateVersion": 12,
  "sites": [...]
}

```

- Using any of the `PATCH` API operations to make a specific change to one of the objects in the schema:

```

PATCH /v1/schemas/601acfed38000070a4ee9ec0?enableVersionCheck=true

[
  {
    "op": "replace",
    "path": "/templates/Template1/displayName",
    "value": "new_name",
    "_updateVersion": 12
  }
]

```

When the request is made, the API will increment the current schema version by 1 (from 12 to 13) and attempt to create the new version of the schema. If the new version does not yet exist, the operation will succeed and the schema will be updated; if another API call (with `enableVersionCheck` enabled) or the UI have modified the schema in the meantime, the operation fails and the API call will return the following response:

```

{
  "code": 400,
  "message": "Update failed, object version in the DB has changed, refresh your client and retry"
}

```

Creating Schemas and Templates

Before you begin

- You must have at least one available tenant that you want to incorporate into your site.

For more information, refer to [Tenants and Tenant Policies, on page 5](#).

Procedure

Step 1 Log in to your Nexus Dashboard and open the Nexus Dashboard Orchestrator service.

Step 2 Create a new schema.

- a) From the left navigation pane, choose **Application Management > Schemas**.
- b) On the Schemas page, click **Add Schema**.
- c) In the schema creation dialog, provide the **Name** and optional description for the schema and click **Add**.

By default, the new schema is empty, so you need to add one or more templates.

Step 3 Create a template.

- a) In the schema page, click **Add New Template**.
- b) In the **Select a Template type** window, choose **ACI Multi-Cloud** and click **Add**.
 - **ACI Multi-Cloud**—Templates used for Cisco ACI on-premises and cloud sites, which allow template and object stretching between multiple sites. This template supports two deployment types:
 - **Multi-Site** - The template can be associated to a single site (site-local policies) or to multiple sites (stretched policies) and the option should be selected for Multi-Site Network (ISN) or VXLAN intersite communication to allow template and object stretching between multiple sites.
 - **Autonomous** - The template can be associated to one or more sites that are operated independently and are not connected through an Inter-Site Network (no intersite VXLAN communication).

Because autonomous sites are by definition isolated and do not have any intersite connectivity, there is no shadow object configuration across sites and no cross-programming of ptags or VNIDs in the spine switches for intersite traffic flow.

The autonomous templates also allow for significantly higher deployment scale.

The following sections focus primarily on this type of templates.

- **NDFC**—Templates designed for Cisco Nexus Dashboard Fabric Controller (formerly Data Center Network Manager) sites.

This guide described Nexus Dashboard Orchestrator configurations for on-premises Cisco ACI fabrics. For information on working with Cisco NDFC sites, see the [Cisco Nexus Dashboard Orchestrator Configuration Guide for NDFC Fabrics](#) instead.

- **Cloud Local**—Templates designed for specific Cloud Network Controller use cases, such as Google Cloud site connectivity, and cannot be stretched between multiple sites.

This guide described Nexus Dashboard Orchestrator configurations for on-premises Cisco ACI fabrics. For information on working with Cloud Network Controller fabrics, see the Nexus Dashboard Orchestrator [use case library](#) instead.

- c) In the right sidebar, provide the **Display Name** for the template.
- d) (Optional) Provide a **Description**.
- e) From the **Select a Tenant** dropdown, select the Tenant for this template.

Keep in mind, the user account you're using to create a new schema must be associated with the tenant you are trying to add to it, otherwise the tenant will not be available in the dropdown menu. Associating a user account with a tenant is described in [Tenants and Tenant Policies, on page 5](#).

- f) In the template view page, click **Save**.

You must save the template after this initial configuration for additional options (such as site association) to become available.

- g) Repeat this step to create any additional templates.

For more information on schema and template design, see [Schema and Template Design Considerations, on page 17](#).

Step 4 Assign the templates to sites.

You deploy fabric configuration by deploying one template at a time to one or more sites. So you need to associate the template with at least one site where you want to deploy the configuration.

- a) In the template view page, click **Actions** and choose **Sites Association**.
- b) In the **Add Sites to <template>** dialog, select one or more sites where you want to deploy the template and click **Ok**.

What to do next

After you have created a schema and one or more templates, you can proceed with editing the templates as described in the following sections of this document based on your specific use cases. After you finish defining configurations, you can deploy the templates as described in [Deploying Templates, on page 51](#).

Importing Schema Elements From APIC Sites

You can create new objects and push them out to one or more sites or you can import existing site-local objects and manage them using the Multi-Site Orchestrator. This section describes how to import one or more existing objects, while creating new objects is described later on in this document.

When importing policies from APIC into NDO, the common practice is to import some objects, such as VRFs or contracts, into a stretched template and other objects, such as non-stretched EPGs or BDs, into site-local templates.

Prior to Release 3.1(1), importing an object into a site-local template that referenced another object that is part of a stretched template presented certain challenges, for example:

- If a referenced object already exists in NDO and a new object is imported with the **Include Relations** option enabled, NDO would throw an error when trying to deploy the site-local template because of object duplication since the referenced object already existed.

- However, not importing the referenced object (**Include Relations** option disabled) would require an administrator to perform manual mapping with the referenced object after the import.

Beginning with Release 3.1(1), when importing an object into a site-local template that has references with another object that is part of a different template (in the same or a different schema), the references are automatically resolved by NDO. In such cases, the **Import Relations** option will be grayed-out in the UI for the object that is being imported and a warning tooltip will provide additional info, such as: *[Referenced Object] already exists in [Template]. Existing relations are imported by default.* While such objects are imported with their relations by default, you can change the references once the import operation is completed, for example by re-mapping a BD to a different VRF. The new behavior applies to all configuration objects that can be imported.

To import one or more objects from sites:

Procedure

-
- Step 1** Open the **Schema** where you want to import objects.
 - Step 2** In the left sidebar, select the **Template** where you want to import objects.
 - Step 3** In the main pane click the **Import** button and select the **Site** from which you want to import.
 - Step 4** In the **Import from <site-name>** window that opens, select one or more objects.

Note

The names of the objects imported into NDO must be unique across all sites. Importing different objects with duplicate names will cause a schema validation error and the import to fail. If you want to import objects that have the same name, you must first rename them.

- Step 5** (Optional) Enable the **Import Relations** knob to import all related objects.

For example, when importing a BD, enabling the **Import Relations** knob will import the associated VRF as well.

Note

As described previously, the **Import Relations** knob will be enabled by default and cannot be disabled for objects whose related objects already exist in NDO.

Configuring VRFs

This section describes how to configure a VRF.

Before you begin

You must have the schema and template created and a tenant assigned to the template, as described in [Creating Schemas and Templates, on page 29](#).

Procedure

-
- Step 1** Select the schema and template where you want to create the VRF.

- a) In the main pane, select **+Create Object > VRF**.

Alternatively, you can scroll down to the **VRFs** area, mouse over the tile, and click **Add VRF**.

- b) In the right pane, provide the **Display Name** for the VRF.
c) (Optional) Provide a **Description**.

Step 2 Configure the **On-Premises Properties** for the VRF.

- a) Specify **Policy Control Enforcement Preference**.

Note that you cannot change the Policy Control Enforcement for newly created VRFs and the setting is locked to the `enforced` mode.

However, you can use this to transition any VRF that you import from an APIC site that is configured as `unenforced` to the `enforced` mode after importing it. A typical use case is for brown field deployments where existing VRFs must be converted to `enforced` mode to support stretching them between sites. Once you have transitioned an imported VRF from `unenforced` to `enforced` in NDO, you will not be able to make further changes to this field.

- `Enforced`—Security rules (contracts) will be enforced.
- `Unenforced`—Security rules (contracts) will not be enforced.

- b) (Optional) Enable **IP Data-Plane Learning**.

Defines if IP addresses are learned through data-plane packets for the VRF.

When disabled, IP addresses are not learned from the data-plane packets. Local and remote MAC addresses are still learned, but local IP addresses are not learned from data packets.

Regardless of whether this parameter is enabled or disabled, local IP addresses can still be learned from ARP, GARP, and ND.

- c) (Optional) Enable **L3 Multicast** for the VRF.

For additional information, see [Layer 3 Multicast, on page 251](#).

- d) (Optional) Enable **vzAny** for the VRF.

For additional information, see [vzAny Contracts, on page 313](#).

- e) (Optional) Enable **Preferred Group** for the VRF.

For additional information, see [EPG Preferred Groups Overview and Limitations, on page 201](#)

Configuring Bridge Domains

This section describes how to configure a Bridge Domain (BD).

Before you begin

- You must have the schema and template created and a tenant assigned to the template, as described in [Creating Schemas and Templates, on page 29](#).
- You must have the VRF created as described in [Configuring VRFs, on page 31](#)

Procedure

Step 1 Select the schema and template where you want to create the bridge domain.

Step 2 Create a bridge domain.

- a) In the main pane, select **+Create Object > Bridge Domain**.

Alternatively, you can scroll down to the **Bridge Domains** area, mouse over the tile, and click **Add Bridge Domain**.

- b) In the right pane, provide the **Display Name** for the bridge domain.

- c) (Optional) Provide a **Description**.

Step 3 Configure **On-Premises Properties**.

- a) From the **Virtual Routing & Forwarding** dropdown, select the VRF for this BD.

- b) (Optional) Enable **L2 Stretch**.

- c) (Optional) Enable **Intersite BUM Traffic Allow**.

This option becomes available if you enabled **L2 Stretch**.

- d) (Optional) Enable **Optimized WAN Bandwidth**.

This option becomes available if you enabled **L2 Stretch**.

- e) (Optional) Enable **Unicast Routing**.

If this setting is enabled and a subnet address is configured, the fabric provides the default gateway function and routes the traffic. Enabling unicast routing also instructs the mapping database to learn the endpoint IP-to-VTEP mapping for this bridge domain. The IP learning is not dependent upon having a subnet configured under the bridge domain.

- f) (Optional) Enable **L3 Multicast** for the BD.

For additional information about Layer 3 multicast, see [Layer 3 Multicast, on page 251](#).

- g) (Optional) Choose **L2 Unknown Unicast** mode.

By default, unicast traffic is flooded to all Layer 2 ports. If enabled, unicast traffic flooding is blocked at a specific port, only permitting egress traffic with MAC addresses that are known to exist on the port. The method can be `Flood` or `Hardware Proxy`.

When the BD has L2 Unknown Unicast set to Flood, if an endpoint is deleted the system deletes it from both the local leaf switches as well as the remote leaf switches where the BD is deployed, by selecting Clear Remote MAC Entries. Without this feature, the remote leaf continues to have this endpoint learned until the timer expires.

Note

Modifying the L2 Unknown Unicast setting causes traffic to bounce (go down and up) on interfaces to devices attached to EPGs associated with this bridge domain.

- h) (Optional) Choose **Unknown Multicast Flooding** mode.

This is applicable for IPv4 unknown multicast traffic and is the node forwarding parameter for Layer 3 unknown multicast destinations.

- `Flood` (default)—Unknown IPv4 multicast traffic is flooded on all front panel ports attached with the EPGs associated with this bridge domain. Flooding is not restricted to only M-Router ports of the bridge domain.
- `Optimized Flood`—Send the data only to M-router ports in the bridge domain.

- i) (Optional) Choose **IPv6 Unknown Multicast Flooding** mode.

This is applicable for IPv6 unknown multicast traffic and is the node forwarding parameter for Layer 3 unknown multicast destinations.

- **Flood (default)**—Unknown IPv6 multicast traffic is flooded on all front panel ports attached with the EPGs associated with this bridge domain. Flooding is not restricted to only M-Router ports of the bridge domain.
- **Optimized Flood**—Send the data only to M-router ports in the bridge domain.

- j) (Optional) Choose **Multi-Destination Flooding** mode.

The multiple destination forwarding method for Layer 2 multicast and broadcast traffic.

- **Flood in BD**—Sends the data to all ports on the same bridge domain.
- **Drop**—Drops Packet. Never sends the data to any other ports.
- **Flood in Encapsulation**—Send the data to all the EPG ports with the same VLAN within the bridge domain, except for the protocol packets which are flooded to the entire bridge domain.

Note

This mode is supported only when the **L2 Stretch** option is disabled and is not supported for BDs that are stretched across sites.

- k) (Optional) Enable **ARP Flooding**.

Enables ARP flooding, so that the Layer 2 broadcast domain maps IP addresses to the MAC addresses. If flooding is disabled, unicast routing will be performed on the target IP address.

Enables ARP flooding, so that ARP request will be flooded inside the Layer 2 broadcast domain. If the BD is stretched across sites, enabling ARP flooding is only possible in conjunction with enabling **Intersite BUM Traffic Allow**. When ARP flooding is disabled, the leaf receiving the ARP request from a locally connected endpoint will forward it directly to the remote leaf where the target endpoint of the ARP request is connected (if the IP for the remote endpoint is known in the endpoint table) or to the spines (if the IP for the remote endpoint is not known in the endpoint table).

If you set the **L2 Unknown Unicast** mode to **Flood**, the **ARP Flooding** cannot be disabled. If the **L2 Unknown Unicast** mode is set to **Hardware Proxy**, ARP flooding can be enabled or disabled.

- l) (Optional) Provide **Virtual MAC Address**.

The BD virtual MAC address and the subnet virtual IP address must be the same for all ACI fabrics for that bridge domain. Multiple bridge domains can be configured to communicate across connected ACI fabrics. The virtual MAC address and the virtual IP address can be shared across bridge domains.

Note

Virtual MAC along with virtual IP subnet should be used only for migration of individual sites to NDO-orchestrated multi-site fabric. Once the migration is completed, these flags can be disabled.

Step 4 Add one or more **Subnets** for the BD.

- a) Click **+Add Subnet**.

An **Add New Subnet** window opens.

- b) Enter the subnet's **Gateway IP** address and a **Description** for the subnet you want to add.
- c) If necessary, enable **Treat as virtual IP address** option.

This option along with the **Virtual MAC Address** on the BD can be used for migration scenarios from individual Common Pervasive Gateway configuration to NDO-managed Multi-Site deployments.

- d) Select the **Scope** for the subnet.

The network visibility of the subnet.

- **Private to VRF**—Prevents the subnet from being announced via L3Out toward an external network domain.
- **Advertised Externally**—The subnet can be announced via L3Out toward an external network domain.

- e) (Optional) Enable **Shared Between VRFs**.

Shared between VRFs—The subnet can be shared with and exported to multiple contexts (VRFs) in the same tenant or across tenants as part of a shared service. An example of a shared service is a routed connection to an EPG present in another context (VRF) in a different tenant. This enables traffic to pass in both directions across contexts (VRFs). An EPG that provides a shared service must have its subnet configured under that EPG (not under a bridge domain), and its scope must be set to advertised externally, and shared between VRFs.

Shared subnets must be unique across the contexts (VRF) involved in the communication. When a subnet under an EPG provides a Layer 3 external network shared service, such a subnet must be globally unique within the entire ACI fabric.

- f) Leave the **No Default SVI Gateway** option unchecked.

Enabling this options means that only the proxy route (subnet route to spine proxy) is programmed on the leaf switches and no SVI is created, which means SVI cannot be used as the gateway.

We recommend that SVI is created by the BD subnet as the gateway and the **No Default SVI Gateway** option is enabled on the EPG instead because EPG subnets should only be used for route leaking.

- g) (Optional) Enable **Querier** option.

Enables IGMP Snooping on the subnet

- h) (Optional) Enable **Primary** option to designate the subnet as primary.

There can be one primary IPv4 subnet and one primary IPv6 subnet.

- i) Click **Save**.

Step 5 (Optional) Add an **IGMP Interface Policy**.

Beginning with Release 4.0(1), you can configure a number of Tenant Policy templates and associate them with policy objects. For more information, see [Creating Tenant Policy Templates, on page 8](#).

Step 6 (Optional) Add an **IGMP Snoop Policy**.

Beginning with Release 4.0(1), you can configure a number of Tenant Policy templates and associate them with policy objects. For more information, see [Creating Tenant Policy Templates, on page 8](#).

Step 7 (Optional) Add an **MLD Snoop Policy**.

Beginning with Release 4.0(1), you can configure a number of Tenant Policy templates and associate them with policy objects. For more information, see [Creating Tenant Policy Templates, on page 8](#).

Step 8 (Optional) Add a **DHCP Policy**.

For additional information, see [DHCP Relay, on page 193](#).

Step 9 Configure the bridge domain's site-local properties as necessary.

In addition to the template-level configurations, you can also define one or more site-local properties for the bridge domain, as described in [Configuring Bridge Domain's Site-Local Properties, on page 36](#)

Configuring Bridge Domain's Site-Local Properties

In addition to the template-level properties you typically configure for the object when you create it in a template, you can also define one or more properties that are specific to each site to which you assign the template.

When you deploy the object to more than 1 site, the same template-level configurations are deployed to all sites, while the site-local configurations are deployed to those specific sites only.

Before you begin

You must have:

- Created the bridge domain and configured its template-level properties, as described in [Configuring Bridge Domains, on page 32](#).
- Assigned the template that contains the bridge domain to one or more sites.

Procedure

- Step 1** Open the schema that contains the template with the bridge domain.
- Step 2** In the left sidebar, select the template that contains the bridge domain under the specific site that you want to configure.
- Step 3** In the main pane, select the bridge domain.
- For most fields, you will see the values you have configured at the template level, which you cannot edit here.
- Step 4** Click **+L3Out** to add an L3Out.
- This is required to advertise the BD subnet out of the remote L3Out and ensure that inbound traffic to the BD can be maintained even if the local L3Out failed. In this case, you would also need to configure the subnet with the `Advertised Externally` flag. For more information, see the [Intersite L3Out, on page 205](#) use case.
- Step 5** Enable **Host Route**.
- This enables Host Based Routing on the bridge domain. When this knob is enabled, the border leaf switches will also advertise individual end-point (EP) host-routes (/32 or /128 prefixes) along with the subnet. The host-route information is advertised only if the host is connected to the local Pod. If the EP is moved away from the local Pod or once the EP is removed from EP database, the route advertisement is then withdrawn.
- Step 6** If necessary, change the **SVI MAC Address**.
- The SVI MAC addresses must be unique per site, when virtual MAC and virtual IP are enabled for Common Pervasive Gateway (CPG) scenario. This field can also be used when CPG is not enabled, which will change the default router MAC of the BD
- Step 7** Add one or more **Subnets** for the BD.
- The concept is the same as adding subnets to the BD at the template level, except the subnets will be configured for the bridge domain on this specific site only.

- a) Click **+Add Subnet**.

An **Add New Subnet** window opens.

- b) Enter the subnet's **Gateway IP** address and a **Description** for the subnet you want to add.
c) Select the **Scope** for the subnet.

The network visibility of the subnet.

- **Private to VRF**—The subnet applies only within its tenant.
- **Advertised Externally**—The subnet can be exported to a routed connection.

- d) (Optional) Enable **Shared Between VRFs**.

Shared between VRFs—The subnet can be shared with and exported to multiple contexts (VRFs) in the same tenant or across tenants as part of a shared service. An example of a shared service is a routed connection to an EPG present in another context (VRF) in a different tenant. This enables traffic to pass in both directions across contexts (VRFs). An EPG that provides a shared service must have its subnet configured under that EPG (not under a bridge domain), and its scope must be set to advertised externally, and shared between VRFs.

Shared subnets must be unique across the contexts (VRF) involved in the communication. When a subnet under an EPG provides a Layer 3 external network shared service, such a subnet must be globally unique within the entire ACI fabric.

- e) (Optional) Enable **No Default SVI Gateway**.

Enabling this options means that only the proxy route (subnet route to spine proxy) is programmed on the leaf switches and no SVI is created, which means SVI cannot be used as the gateway.

We recommend that SVI is created by the BD subnet as the gateway and the **No Default SVI Gateway** option is enabled on the EPG instead because EPG subnets should only be used for route leaking.

- f) (Optional) Enable **Querier**.

Enables IGMP Snooping on the subnet

- g) (Optional) Enable **Primary** option to designate the subnet as primary.

There can be one primary IPv4 subnet and one primary IPv6 subnet.

- h) Click **Save**.

Configuring Application Profiles and EPGs

This section describes how to configure an Application Profile and an EPG.

Before you begin

You must have the schema and template created and a tenant assigned to the template, as described in [Creating Schemas and Templates, on page 29](#).

This section also assume you have a Contract and a Bridge Domain created.

Procedure

-
- Step 1** Select the schema and template where you want to create the application profile.
- Step 2** Create an application profile.
- a) In the main pane, select **+Create Object > Application Profile**.
Alternatively, you can scroll down to the **Application Profile** area, mouse over the tile, and click **Add Application Profile**.
 - b) In the right pane, provide the **Display Name** for the application profile.
You can create application profiles with the same name in different templates without any conflicts. You cannot however create other objects (such as VRFs, BDs, EPGs) with the same name in different templates if they will be deployed to the same site and tenant.
 - c) (Optional) Provide a **Description**.
- Step 3** Create an EPG.
- a) In the main pane, select **+Create Object > EPG**, then select the application profile where you want to create the EPG.
Alternatively, you can scroll down to the specific **Application Profile** area, mouse over the **EPGs** tile, and click **Add EPG**.
 - b) In the right pane, provide the **Display Name** for the EPG.
 - c) (Optional) Provide a **Description**.
- Step 4** Add a contract for the EPG.
- Creating contracts and filters is described in detail in [Configuring Contracts and Filters, on page 42](#). If you already have a contract created:
- a) Click **+ Contract**.
 - b) On the **Add Contract** dialog, enter the contract name and type.
 - c) Click **SAVE**.
- Step 5** From the **Bridge Domain** dropdown, select the bridge domain for this EPG.
- If you are configuring an on-premises EPG, you must associate it with a bridge domain.
- Step 6** (Optional) Click **+ Subnet** to add a subnet to your EPG.
- You may choose to configure a subnet on the EPG level rather than the bridge domain level, for example for a VRF route-leaking use-case.
- a) On the **Add Subnet** dialog, enter the **Gateway IP** address and a description for the subnet you plan to add.
 - b) In the **Scope** field select either **Private to VRF** or **Advertised Externally**.
 - c) Click the check box for **Shared Between VRFs** if appropriate.
 - d) Click the check box for **No Default SVI Gateway** if appropriate.
 - e) Click **OK**.
- Step 7** (Optional) Enable microsegmentation.
- If you are configuring a microsegmentation EPG (uSeg), you must provide one or more uSeg attributes for matching endpoints to the EPG.

- a) Check the **uSeg EPG** checkbox.
- b) Click **+uSeg Attribute**.
- c) Provide the **Name** and **Type** for the uSeg attribute.
- d) Based on the attribute type you have selected, provide the attribute details.

For example, if you have selected `MAC` for the attribute type, provide the MAC address to identify an endpoint in this EPG.

- e) Click **SAVE**.

Step 8 (Optional) Enable intra-EPG isolation.

By default, endpoints in EPG can freely communicate with each other. If you would like to isolate the endpoints from each other, set the isolation mode to **Enforced**.

Intra-EPG endpoint isolation policies provide full isolation for virtual or physical endpoints; no communication is allowed between endpoints in an EPG that is operating with isolation enforced. Isolation-enforced EPGs reduce the number of EPG encapsulations required when many clients access a common service but are not allowed to communicate with each other.

Step 9 (Optional) Enable Layer 3 multicast for the EPG.

For additional information about Layer 3 multicast, see [Layer 3 Multicast, on page 251](#)

Step 10 (Optional) Enable preferred group membership for the EPG.

The Preferred Group feature allows you to include multiple EPGs within a single VRF to allow full communication between them with no need for contracts to be created. For additional information about EPG preferred group, see [EPG Preferred Groups Overview and Limitations, on page 201](#)

Step 11 Configure the EPG's site-local properties as necessary.

In addition to the template-level configurations, you can also define one or more site-local properties for the EPG, as described in [Configuring EPG's Site-Local Properties, on page 39](#)

Configuring EPG's Site-Local Properties

In addition to the template-level properties you typically configure for the object when you create it in a template, you can also define one or more properties that are specific to each site to which you assign the template.

When you deploy the object to more than 1 site, the same template-level configurations are deployed to all sites, while the site-local configurations are deployed to those specific sites only.

Before you begin

You must have:

- Created the application profile and EPG and configured the template-level properties, as described in [Configuring Application Profiles and EPGs, on page 37](#).
- Assigned the template that contains the bridge domain to one or more sites.

Procedure

Step 1 Open the schema that contains the template with the EPG.

Step 2 In the left sidebar, select the template that contains the EPG under the specific site that you want to configure.

Step 3 In the main pane, select the EPG.

For most fields, you will see the values you have configured at the template level, which you cannot edit here.

Step 4 Add one or more **Subnets** for the EPG.

a) Click **+Add Subnet**.

An **Add New Subnet** window opens.

b) Enter the subnet's **Gateway IP** address and a description for the subnet you want to add.

c) Select the **Scope** for the subnet.

The network visibility of the subnet.

- **Private to VRF**—Prevents the subnet from being announced via L3Out toward an external network domain.
- **Advertised Externally**—The subnet can be announced via L3Out toward an external network domain.

d) (Optional) Enable **Shared Between VRFs**.

Shared between VRFs—The subnet can be shared with and exported to multiple contexts (VRFs) in the same tenant or across tenants as part of a shared service. An example of a shared service is a routed connection to an EPG present in another context (VRF) in a different tenant. This enables traffic to pass in both directions across contexts (VRFs). An EPG that provides a shared service must have its subnet configured under the BD (not under the EPG), and its scope must be set to advertised externally, and shared between VRFs.

Shared subnets must be unique across the contexts (VRF) involved in the communication. When a subnet under an EPG provides a Layer 3 external network shared service, such a subnet must be globally unique within the entire ACI fabric.

e) (Optional) Enable **No Default SVI Gateway**.

Enabling this options means that only the proxy route (subnet route to spine proxy) is programmed on the leaf switches and no SVI is created, which means SVI cannot be used as the gateway.

We recommend enabling this option on the EPG subnets, which should only be used for route leaking and leaving this option disabled on the BD subnets so that the SVI can be used as a gateway.

f) Click **Ok** to save.

Step 5 Add one or more **Static ports**.

a) Click **+Static Port**.

b) From the **Path Type** dropdown, select the type of port.

c) If configuring a physical interface, select the **Pod**

d) Choose whether you want to configure a single port or a range of ports.

For the interface configuration, you will have an option to do it either by entering a single **Leaf** and a **Path** or by entering a range of **Leaf** eg 120-125 and **Path** eg 1/17-20. You will also have an option to enter a range of **Leaf** and associate it with one single **Path**, or enter a range of **Path** for one single **Leaf**.

However, after the configuration it will still be displayed as individual ports in the UI and will require individual changes for any future updates.

e) Select the **Port Encap VLAN**.

When manually configuring the port encap on a domain for an EPG, the VLAN ID must belong to a static VLAN block within a dynamic VLAN pool.

If EPG is enabled for microsegmentation at the template level, when a **Primary MICRO-SEG VLAN** is configured, the **Port Encap VLAN** is configured as an Isolated Secondary VLAN for the Primary VLAN. Traffic is sent from the host to the leaf using the secondary VLAN and return traffic from the leaf to the host is sent using the primary VLAN.

f) (Optional) Select the **Primary MICRO-SEG VLAN**.

The VLAN identifier for microsegmentation

g) (Optional) Select the **Deployment Immediacy**.

Once policies are downloaded to the leaf nodes, deployment immediacy can specify when the policy is pushed into the hardware policy CAM:

- **Immediate**—Specifies that the policy is programmed in the hardware policy CAM as soon as the policy is downloaded in the leaf software.
- **On Demand**—Specifies that the policy is programmed in the hardware policy CAM only when the first packet is received through the data path. This process helps to optimize the hardware space.

h) (Optional) Select the **Mode**.

The mode of the static association with the path. EPG tagging refers to configuring a static path under an EPG:

- **Trunk**—The default deployment mode. Select this mode if the traffic from the host is tagged with a VLAN ID.
- **Access (802.1P)**—Select this mode if the traffic from the host is tagged with a 802.1P tag. When an access port is configured with a single EPG in native 802.1p mode, its packets exit that port untagged. When an access port is configured with multiple EPGs, one in native 802.1p mode, and some with VLAN tags, all packets exiting that access port are tagged VLAN 0 for EPG configured in native 802.1p mode and for all other EPGs packets exit with their respective VLAN tags. Note that only one native 802.1p EPG is allowed per access port.
- **Access (Untagged)**—Select this mode if the traffic from the host is untagged (without VLAN ID). When a leaf switch is configured for an EPG to be untagged, for every port this EPG uses, the packets will exit the switch untagged. Note that when an EPG is deployed as untagged, do not deploy that EPG as tagged on other ports of the same switch.

Step 6 Add one or more **Static Leaf** nodes.

- a) Click **+Static Leaf**.
- b) From the **Leaf** dropdown, select the leaf node you want to add.
- c) (Optional) In the **VLAN** field, provide the VLAN ID for tagged traffic.

Step 7 Add one or more **Domains** nodes.

- a) Click **+Domain**.
- b) Select the **Domain Association Type**.

This is the type of the domain you are adding:

- VMM

- Fibre Channel
- L2 External
- L3 External
- Physical

- c) Select the **Domain Profile** name.
- d) Select the **Deployment Immediacy**.

Deployment immediacy can specify when the policy is pushed:

- **Immediate**—Specifies that the policy is programmed in the hardware policy CAM as soon as the policy is downloaded in the leaf software.
- **On Demand**—Specifies that the policy is programmed in the hardware policy CAM only when the first packet is received through the data path. This process helps to optimize the hardware space.

- e) Select the **Resolution Immediacy**.

Specifies whether policies are resolved immediately or when needed. The options are:

- **Immediate**—Specifies that EPG policies are pushed to the leaf switch nodes upon hypervisor attachment to the VMware vSphere Distributed Switch (VDS). LLDP or OpFlex permissions are used to resolve the hypervisor to leaf node attachments.
- **On Demand**—Specifies that EPG policies are pushed to the leaf switch nodes only when a hypervisor is attached to VDS and a VM is placed in the port group (EPG).
- **Pre-provision**—Specifies that EPG policies are pushed to the leaf switch nodes even before a hypervisor is attached to the VDS. The download pre-provisions the configuration on the switch.

Configuring Contracts and Filters

This section describes how to configure a contract, a filter, and assign the filter to the contract. A filter is similar to an Access Control List (ACL), it is used to filter traffic through contracts associated to EPGs.

Procedure

- Step 1** Select the schema and template where you want to create contract and filter.

You can create the contract in the same or different template as the objects (EPGs and external EPGs) to which you will apply it. If the objects that will use the contract are deployed to different sites, we recommend defining the contract in a template associated to multiple sites. However, this is not strictly required and even if the contract and filters are defined only as local objects in Site1, NDO will create those objects in a remote Site2 when a local EPG or external EPG in Site2 needs to consume or provide that contract.

- Step 2** Create a filter.
 - a) In the main pane, select **+Create Object > Filter**.

Alternatively, you can scroll down to the **Filters** area, mouse over the tile, and click **Add Filter**.

- b) In the right pane, provide the **Display Name** for the filter.
- c) (Optional) Provide a **Description**.

Step 3

Create a filter entry.

- a) In the right pane, click + **Add Entry**.

The filter entry is a combination of network traffic classification properties. You can specify one or more options as described in the following step.

- b) Provide the **Name** for the filter.
- c) Choose the **Ether Type**.

For example, `ip`.

- d) Choose the **IP Protocol**.

For example, `icmp`.

- e) Choose the **Destination Port Range From** and **Destination Port Range To**.

The start and end of the destination ports range. You can define a single port by specifying the same value in both fields or you can define a range of ports from 0 to 65535. You can also choose to specify one of the server types instead of specific port numbers, for example `http`.

- f) Enable **Match only fragments** option.

When enabled, the rule applies to any IP fragment with an offset that is greater than 0 (all IP fragments except the first). When disabled, the rule will not apply to IP fragments with an offset greater than 0 because TCP/UDP port information can only be checked in initial fragments.

- g) Enable **Stateful** option.

When this option is enabled, any traffic coming from the provider back to the consumer will always have to have the `ACK` bit set in the packet or else the packets will be dropped.

- h) Specify **ARP flag** (Address Resolution Protocol).

The **ARP Flag** is used when creating a specific filter for ARP and allows you to specify ARP request or ARP reply.

- i) Choose the **Source Port Range From** and **Source Port Range To**.

The start and end of the source ports range. You can define a single port by specifying the same value in both fields or you can define a range of ports from 0 to 65535. You can also choose to specify one of the server types instead of specific port numbers, for example `http`.

- j) Specify **TCP session rules**.

TCP session rules are used when creating a filter for TCP traffic and allow you to configure `stateful` ACL behavior.

- k) Click **Ok** to save the filter.
- l) Repeat this step to create any additional filter entries for this filter.

You can create and assign multiple filter entries for each filter.

Step 4

Create a contract

- a) In the main pane, select +**Create Object** > **Contract**.

Alternatively, you can scroll down to the **Contract** area, mouse over the tile, and click **Add Contract**.

- b) In the right pane, provide the **Display Name** for the contract.
- c) (Optional) Provide a **Description**.
- d) Select the appropriate **Scope** for the contract.

Contract scope limits the contract's accessibility; the contract will not be applied to any consumer EPG outside the scope of the provider EPG:

- Application Profile
- VRF
- Tenant
- Global

- e) Toggle the **Apply both directions** knob if you want the same filter to apply for both consumer-to-provider and provider-to-consumer directions.

If you enable this option, you will need to provide the filters only once and they will apply for traffic in both directions. If you leave this option disabled, you will need to provide two sets of filter chains, one for each direction.

Note

If you create and deploy a contract with **Apply both directions** enabled, you cannot simply disable the option and re-deploy for the change to apply. To disable this option on an already deployed contract, you must delete the contract, deploy the template, then re-create the contract with the option disabled to correctly change the setting in your fabrics.

- f) (Optional) From the **Service Graph** dropdown, select a service graph for this contract.
- g) (Optional) From the **QoS Level** dropdown, select a value for this contract.

This value specifies the ACI QoS Level that will be assigned to the traffic using this contract. For more information, see [QoS Preservation Across IPN, on page 261](#).

If you leave this at `Unspecified`, the default QoS Level 3 is applied to the traffic.

Step 5

Assign the filters to the contract.

- a) In the main pane for template, select a contract. In the right pane, scroll down to the **Filter Chain** area and click + **Add Filter** to add a filter to the contract.
- b) In the **Add Filter Chain** window that opens, select the filter you added in previous step from the **Name** dropdown menu.
- c) Select the **Action** for the filter.

When adding filters, you can choose whether to permit or deny traffic that matches the filter criteria. For `deny` filters, you can set the priority of the filter to one of four levels: `default`, `low`, `medium`, or `high`; the `permit` filters always have the default priority. For more information on ACI contracts and filters, see [Cisco ACI Contract Guide](#).

- d) Click **Ok** to add the filter to the contract.
- e) If you disabled the `Apply both directions` option on the contract, repeat this step for the other filter chain.
- f) (Optional) You can create and assign multiple Filters to each Contract.

If you want to create additional filter for the same contract:

- Repeat Step 2 and Step 3 to create another filter along with its filter entries.
- Then repeat this step to assign the new filter to this Contract.

Configuring On-Premises External Connectivity

Cisco ACI allows you to establish connectivity to the networks outside your on-premises ACI fabric through the border leaf switches. This connectivity is defined using two constructs, L3Out and External EPG, which provide the configuration options necessary to define security and route maps.

This section describes how to create an L3Out and external EPG in the Nexus Dashboard Orchestrator GUI. The Orchestrator then creates the objects on the APIC site where you deploy the template. Keep in mind that when creating an L3Out from the Orchestrator, only the L3Out container object is created in the APIC and you must still perform the full L3Out configuration (such as nodes, interfaces, routing protocols, and so on) directly in the site's APIC.

While in most cases the L3Out will be created directly at the APIC level and then associated to an external EPG that you create in the Orchestrator, it may be useful to create both here in order to directly associate the L3Out to a VRF which you have created from the Orchestrator.

Before you begin

- You must have the schema and template created and a tenant assigned to the template, as described in [Creating Schemas and Templates, on page 29](#).
- You must have the VRF for the L3Out created as described in [Configuring VRFs, on page 31](#)

Procedure

Step 1 Navigate to the Schema and Template you want to edit.

Step 2 Create an L3Out.

- a) In the schema edit view, scroll down to the **L3Out** area and click + to add a new L3Out.
- b) In the properties pane on the right, provide a display name for the L3Out.
- c) From the **Virtual Routing & Forwarding** dropdown, select the VRF you created for this.

Note

Once the external EPG and L3Out are deployed to sites, if you want to change the VRF of the L3out or the external EPG, you must first undeploy the existing objects, change the VRF on both the L3Out and the external EPG, and then re-deploy the two objects.

Step 3 Create an external EPG.

- a) In the schema edit view, scroll down to the **External EPG** area and click + to add a new External EPG.
- b) In the properties pane on the right, select **On-Prem** for the site type.

Step 4 Configure external EPG's basic properties.

- a) In the right properties sidebar, provide a display name for the External EPG.
- b) From the **Virtual Routing & Forwarding** dropdown, select the VRF you created.

This must be the same VRF that you associated with the L3Out.

Note

Once the external EPG and L3Out are deployed to sites, if you want to change the VRF of the L3out or the external EPG, you must first undeploy the existing objects, change the VRF on both the L3Out and the external EPG, and then re-deploy the two objects.

- c) Click **+Add Contract** to add a contract for the external EPG to communicate with other EPGs.

If you already have the contracts created, you can assign them now. Otherwise, you can come back to this screen to assign any Contracts you plan to create later.

When assigning Contracts:

- If you are associating a contract with the external EPG as provider, choose contracts only from the tenant associated with the external EPG. Do not choose contracts from other tenants.
- If you are associating the contract to the external EPG as consumer, you can choose any available contract.

Step 5 If configuring an external EPG for an on-premises fabric, set **Site Type** to `on-prem` and configure external EPG's on-premises properties.

- a) From the **L3Out** dropdown, select the L3Out you created in a previous step.

Note

You can select the L3Out at the template or site-local level. We recommend configuring the L3Out for the external EPG at the site-local level. To do that, select the template in the left sidebar under the site to which it is assigned. Then select the external EPG and associate an L3Out with it.

- b) Click **+Subnet** to add a subnet.

This may be a classification subnet or a subnet used for route-control.

- c) In the **Add Subnet** window, provide the subnet prefix.

- d) Select the required **Route Control** options.

You can choose one or more for the following options:

- **Export Route Control** enables a route map that allows external prefixes matching the specified subnet prefix to get advertised out of the L3Out. These are the prefixes learned from other L3Outs for transit routing use cases.

If you are adding the `0.0.0.0/0` subnet and enable the export route control option, **Aggregate Export** option becomes available. This allows you to advertise all the external prefixes learned from other L3Outs. If you choose to leave this option disabled, only the default route learned from other L3Out will be advertised out of this L3Out.

- **Import Route Control** configures an ingress route map to give you control over what prefixes you want to import from your L3Out into the fabric. The **Import Route Control** is available only when using BGP as the routing protocol on the L3Out.

If you are adding the `0.0.0.0/0` subnet and enable the import route control option, **Aggregate Import** option becomes available. This works similar to the export route control case except for ingress routes.

- **Shared Route Control** is used for shared L3Out use case and allows prefixes learned from the external router to be advertised to the other VRF that will use this L3Out.

If you enable the shared route control option, **Aggregate Shared Routes** option becomes available. Again, this functions similar to previous two aggregate routes options but is available for non-`0.0.0.0/0` subnets.

- e) Select the **External EPG Classification** options.

You should check the **External Subnets for External EPG** option for the configured subnet(s) to be able to map external entities to this specific External EPG. This allows you to apply a security policy (contract) between those external networks and endpoints belonging to EPG(s) defined inside the fabric. Enabling this flag for a `0.0.0.0/0` prefix ensures that all the external destinations are considered part of this External EPG.

If you enable this option, **Shared Security Import** option becomes available, which allows access from the subnet to the endpoints within the fabric for the inter-VRF (shared services) use case.

For both of these options, access is still subject to contract rules.

- Step 6** If configuring an external EPG for a cloud fabric, set **Site Type** to `cloud` and configure external EPG's cloud properties.
- a) From the **Application Profile** dropdown, select the application profile.
 - b) Click **+Add Selector** to add a cloud endpoint selector for the EPG.

- Step 7** (Optional) If you want to include this external EPG in the preferred group, check the **Include in Preferred Group** checkbox.

For more info about EPG Preferred Group, see [EPG Preferred Groups Overview and Limitations, on page 201](#).

Viewing Schemas

After you have created one or more schemas, they are displayed both on the Dashboard and the Schemas page.

You can use the functionality available on these two pages to monitor the usage and the health of your schemas when they are deployed. You can also access and edit specific areas of the implemented schema policies using the Nexus Dashboard Orchestrator GUI.

Bulk Update for Template Objects

The bulk update feature allows you to update multiple properties on multiple different objects of the same type within a template at once. For example, you can enforce Infra EPG Isolation on two or more EPGs at the same time, instead of having to modify each object individually. When using this workflow, all selected objects must be of the same type, for example, you cannot choose to update an EPG and a BD simultaneously.

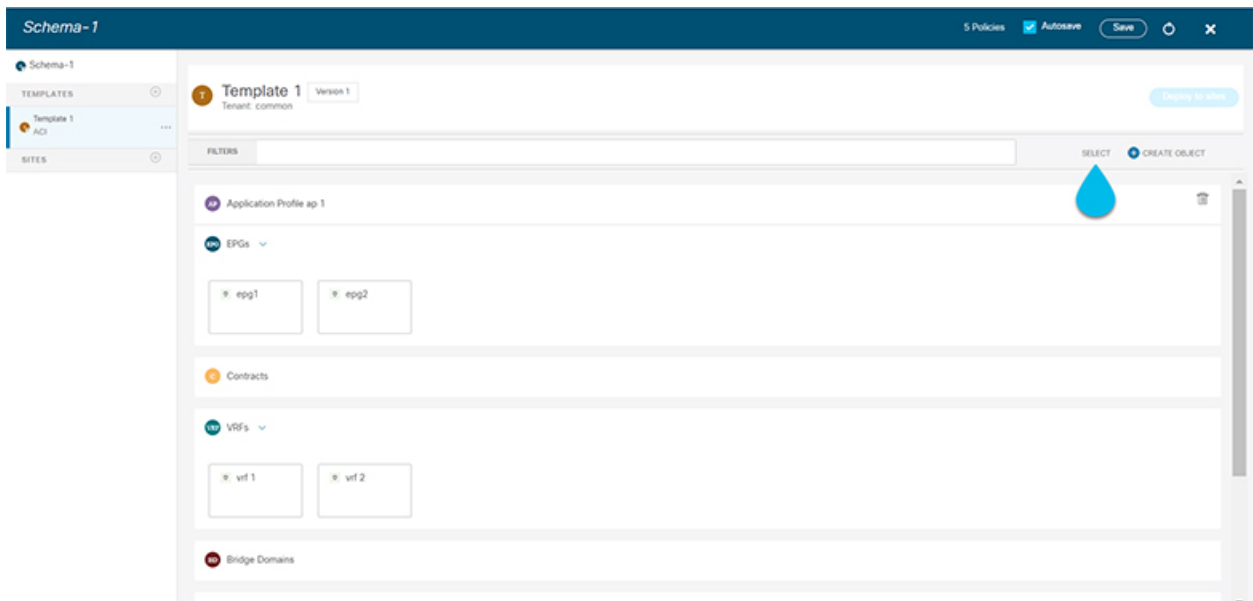
If the selected objects already have different property values configured on them, the update will overwrite those properties with the values you provide. This feature allows you update template-level object properties for on-premises; updating site-local properties and cloud properties are not supported.



Note This feature is supported for Cisco APIC and Cisco NDFC fabrics only; it is not supported for Cisco Cloud Network Controller sites.

Procedure

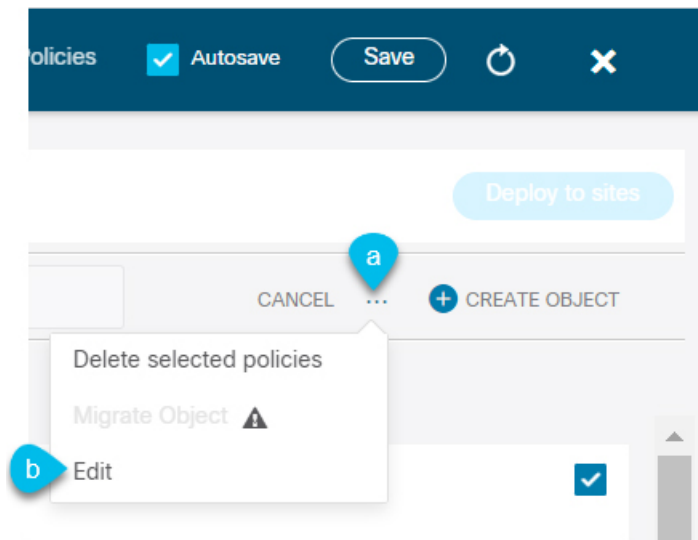
-
- Step 1** Navigate to the schema and template that contains the objects you want to update.
- Step 2** From the main pane, choose **Select**. It will allow you to choose multiple objects of same type.



Step 3 After selecting all the objects that you want to update.

- Choose “...” right next to the cancel option.
- From the dropdown Choose "Edit".

If you choose objects of different type, you won't see the Edit option in the dropdown.

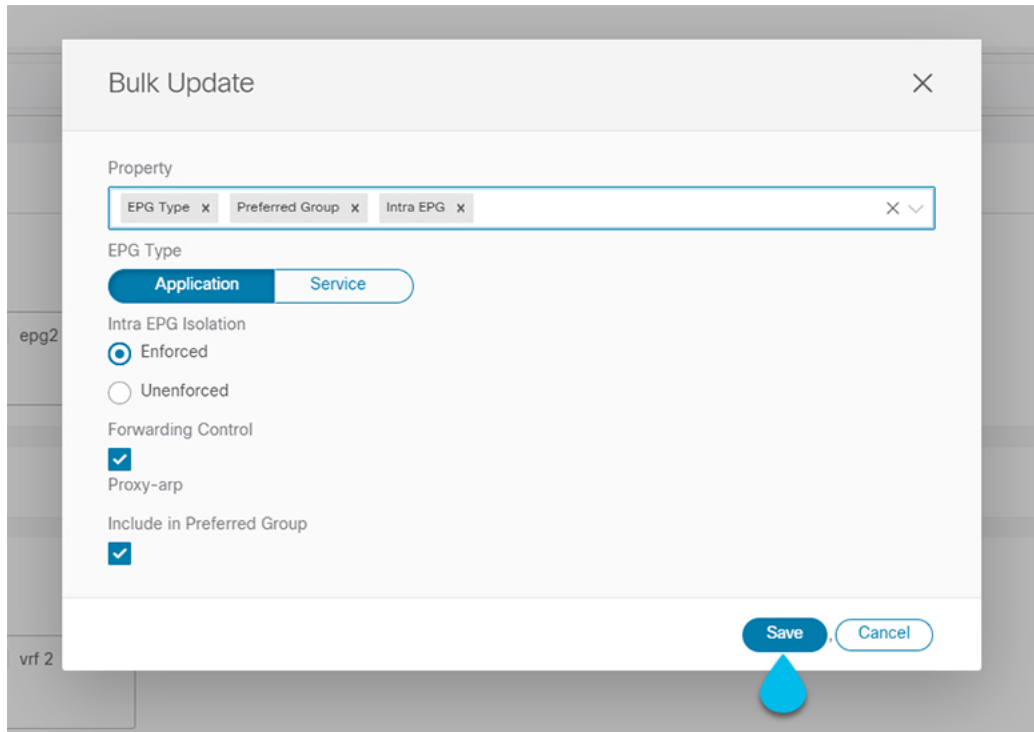


Step 4 After choosing “Edit”, a pop-up will show up. It will show you a subset of the properties for the objects you selected. You can update the following properties based on the type of objects you selected.

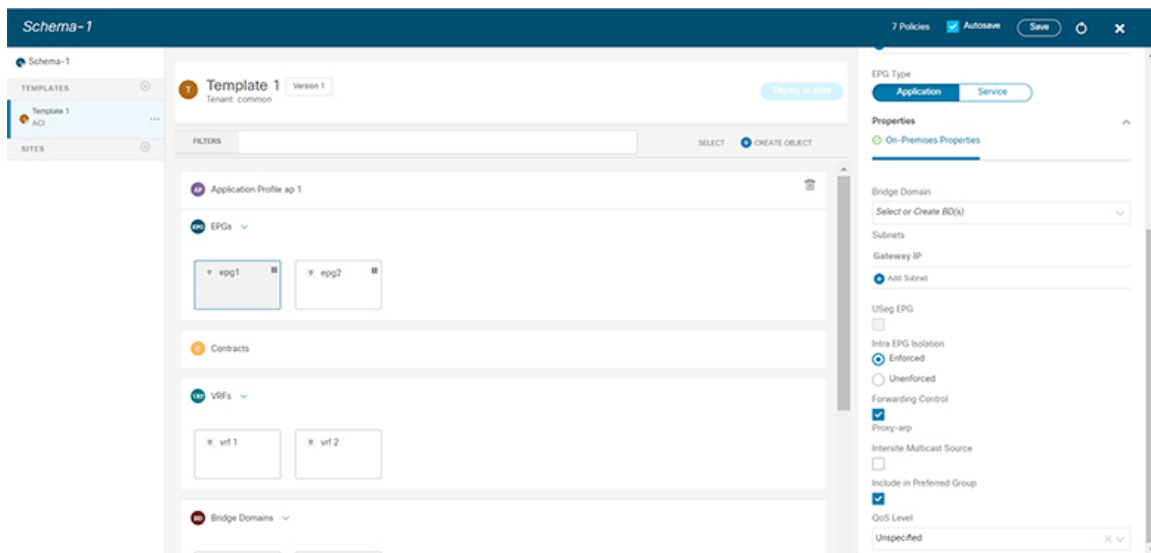
- EPG:** Bridge Domain, Contract, EPG Type, Infra EPG, Preferred Group.
- Contracts:** Scope, Filter Chain, QOS Level.
- VRF:** IP Data-Plane Learning.

- d. **Bridge Domain:** Virtual Routing and Forwarding, L2 Stretch, L2 Unknown Unicast, Unknown Multicast Flooding, IPv6 Unknown Multicast Flooding, Multi Destination Flooding, DHCP Policies, Unicast Routing.
- e. **External EPG:** Contract, External EPG Type, Preferred Group.

Step 5 After selecting all the fields, you wish to update. Choose “Save” which will implement the update you just made.



Step 6 As you save the updates, you can see the changes you've made.



Assigning Templates to Sites

This section describes how to assign a template to sites.

Before you begin

You must have the schema, template, and any objects you want to deploy to sites already created, as described in previous sections of this document.

Procedure

- Step 1** Navigate to the schema that contains one or more templates that you want to deploy.
- Step 2** In the left sidebar, select the template that you want to assign to sites.
- Step 3** In the **Template Properties** view, click **Actions** and choose **Sites Association**.
The **Associate Sites to <template-name>** window opens.
- Step 4** In the **Associate Sites** window, check the checkbox next to the sites where you want to deploy the template.
Note that some sites may not be available for assignment depending on the type of the template you selected and the intersite connectivity between sites:
- If you are assigning a `Cloud Local` template, you will be able to assign it only to a single cloud site.
 - When assigning templates to multiple sites, the intersite connectivity between those sites must be established using BGP-EVPN protocol. If you select a site that has partial mesh connectivity, any site to which there is no intersite connectivity or intersite connectivity is established using BGP-IPv4 will be grayed out and unavailable for assignment.
- Step 5** Click **Ok**.
You deploy one template at a time, so you must associate the template with at least one site before you can deploy it.
-

Disassociating Template from Sites

You can choose to disassociate a template from a site without undeploying it. This allows you to preserve any configuration deployed to the site from NDO while removing the template-site association in the schema. The managed object and policy ownership is transferred from NDO to the site's controller.

Before you begin

- The template and its configuration must already be deployed to a site.
- The template must be deployed to a single site only and not stretched across sites.
- The objects defined in the template must not be deployed as shadow objects in other sites.

Procedure

-
- Step 1** Log in to your Nexus Dashboard Orchestrator GUI.
- Step 2** From the left navigation menu, select **Application Management > Schemas**.
- Step 3** Click the schema that contains the template you want to disassociate.
- Step 4** In the Schema view, select the template under the specific site from which you want to disassociate it.
- Step 5** From the **Actions** menu, select **Disassociate Template**.
- Step 6** In the confirmation window, click **Confirm Action**.
-

Deploying Templates

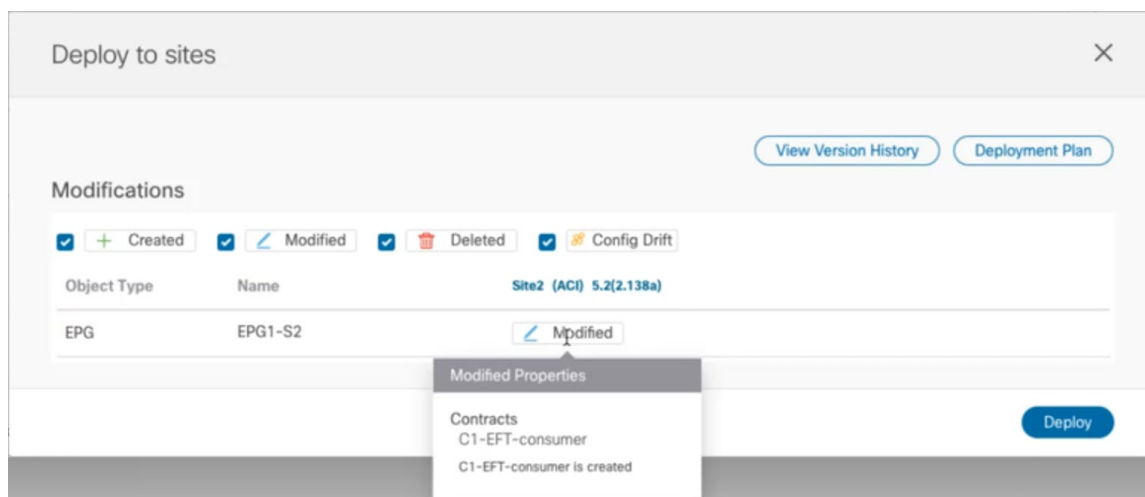
This section describes how to deploy new or updated policies to the ACI fabrics.

Before you begin

- You must have the schema, template, and any objects you want to deploy to sites already created and the templates assigned to one or more sites, as described in previous sections of this document.
- If template review and approval is enabled, the template must also be already approved by the required number of approvers as described in [Template Review and Approval, on page 60](#).
- Ensure that you understand the required deployment order and object dependencies that are described in [Schema and Template Design Considerations, on page 17](#).

Procedure

-
- Step 1** Navigate to the schema that contains the template that you want to deploy.
- Step 2** From the **View** dropdown menu, select the template you want to deploy.
- Step 3** In the template view, click **Deploy to sites**.
- The **Deploy to sites** window opens that shows the summary of the objects to be deployed.
- Step 4** If you have made changes to your template, review the **Deployment Plan** to verify the new configuration.
- If you have previously deployed this template but made no changes to it since, the **Deploy** summary will indicate that there are no changes and you can choose to re-deploy the entire template. In this case, you can skip this step.
- The **Deploy to sites** window will show you a summary of the configuration differences that will be deployed to sites. The following screenshots show a simple example of adding a `consumer` contract to an existing EPG (`EPG1-S2`) in `Site2`.
- Note**
- In this case only the difference in configuration is deployed to the sites. If you want to re-deploy the entire template, you must deploy once to sync the differences and then redeploy again to push the entire configuration as described in the previous paragraph.



You can also filter the view using the **Created**, **Modified**, and **Deleted** checkboxes for informational purposes, but keep in mind that all of the changes are still deployed when you click **Deploy**.

Here you can also choose to:

- **View Version History**, which shows the complete version history and incremental changes made between versions. Additional information about version history is available in [Viewing History and Comparing Previous Versions, on page 57](#).
- Check the **Deployment Plan** to see a visualization and an XML payload of the configuration that will be deployed from this template.

This feature provides better visibility into configuration changes that the Orchestrator will provision to the different fabrics that are part of your Multi-Site domain after you make a change to the template and deploy it to one or more sites.

Unlike earlier releases of the Nexus Dashboard Orchestrator, which still provided a list of specific changes made to the template and site configuration, the Deployment Plan provides full visibility into all the objects that the deployment of the template would provision across the different fabrics. For example, depending on what change you make, shadow objects may be created in multiple sites even if the specific change is applied to only a single site.

Note

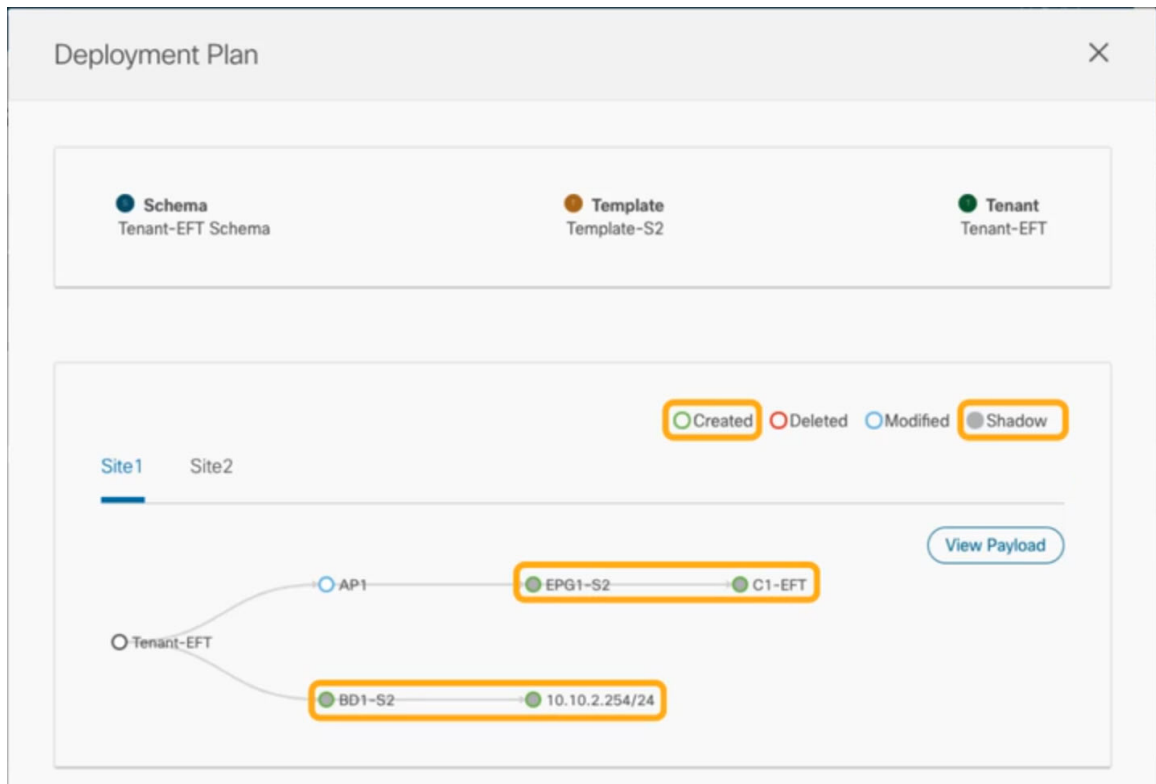
We recommend verifying your changes using the Deployment Plan as described in this step before deploying the template. The visual representation of the configuration changes can help you reduce potential errors from deploying unintended configuration changes.

- Click the **Deployment Plan** button.

Continuing with the same example shown in the previous step, where a consumer contract was added to an existing EPG in `Site2`, the Deployment Plan allows you to also see that there are additional changes to be deployed to `Site1` as a result of the change to `Site2`.

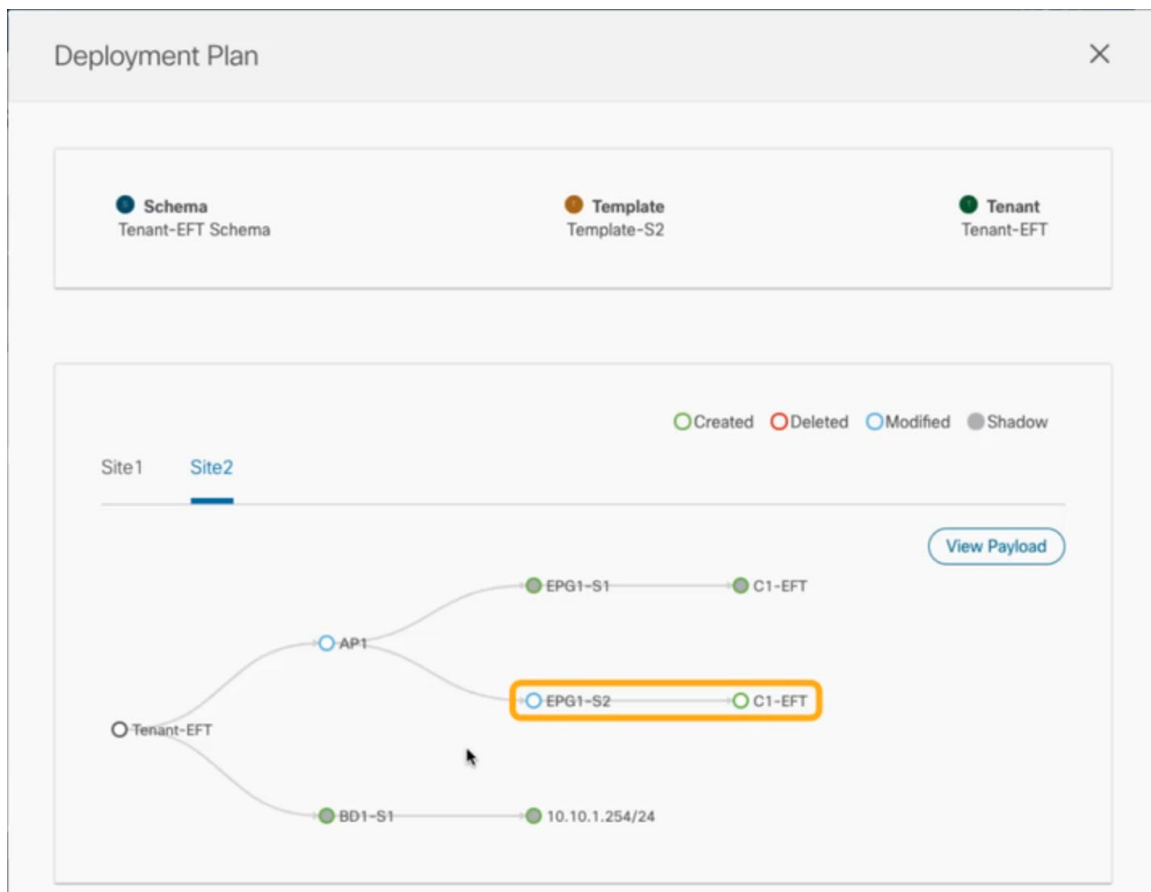
- Verify your changes in the first listed site.

Based on the highlighted legend, you can see that the Orchestrator will create the shadow objects in `Site1` that are required by the contract you added to an EPG in `Site2`.



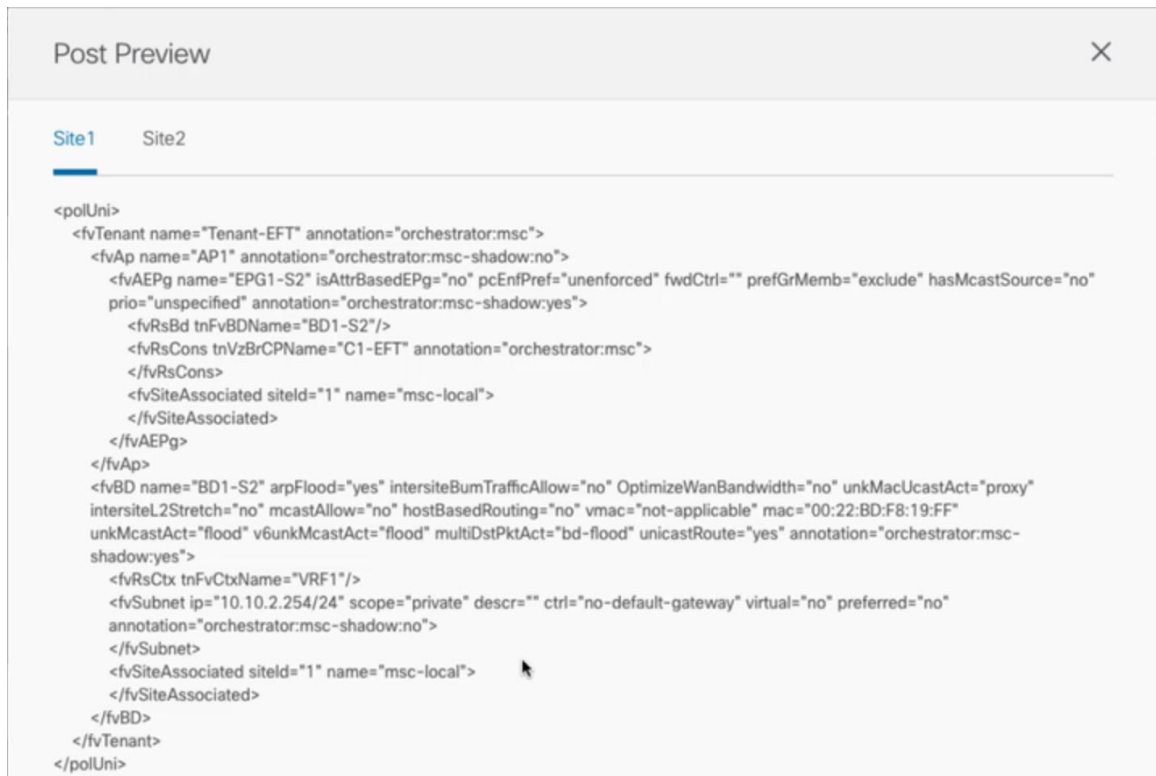
- c) Repeat the previous substep to verify the changes in other sites

Here you can see the change you made explicitly to the EPG (EPG1-S2) in Site2 when you assigned the contract (C1-EFT) to it, as well as the shadow objects for the EPG (EPG1-S1) in the other site, which is providing that contract.



- d) (Optional) Click **View Payload** to see the XML payload for each site.

In addition to the visual representation of the new and modified objects, you can also choose to **View Payload** for the changes in each site:



Step 5 In the **Deploy to sites** window, click **Deploy** to deploy the template.

Undeploying Templates

This section describes how to undeploy a template from a site.

Before you begin

- Ensure that you have not made any changes to the template since you last deployed it.

Undeploying a template that was modified since it was last deployed may create a configuration drift because the set of objects deployed with the template would be different than the set of objects you try to undeploy after making changes to the template.

- If you are undeploying a template that contains VRFs that are used in route leak configurations, the route leaks must be deleted before you can undeploy that template.

Procedure

Step 1 Select the schema that contains the template you want to undeploy.

Step 2 From the **View** dropdown, select the template you want to undeploy.

Step 3 From the **Actions** menu, click **Undeploy from all sites**.

Template Versioning

A new version of the template is created every time it is saved. From within the NDO UI, you can view the history of all configuration changes for any template along with information about who made the changes and when. You can also compare any of the previous versions to the current version.

New versions are created at the template level, not schema level, which allows you to configure, compare, and roll back each template individually.

Template versions are created and maintained according to the following rules:

- All template versions are either `Deployed` or `Intermediate`.
`Deployed`—versions of the template that have been deployed to sites.
`Intermediate`—versions of the template that have been modified and saved, but not deployed to sites.
- A maximum of 20 `Deployed` and 20 `Intermediate` versions per template can be stored at any given time.
- When a new `Intermediate` version is created that would exceed the 20 version limit, the earliest existing `Intermediate` version is deleted.
- When a template is deployed and a new `Deployed` version is created, all `Intermediate` versions are deleted. If the new `Deployed` version exceeds the 20 version limit, the earliest existing `Deployed` version is deleted.
- Tagging a version `Golden` does not affect the number of stored template versions.
- A template that is tagged `Golden` cannot be deleted.
You must untag the template first before you can delete it.
- When a template is modified and saved or deployed, any versions that exceed the 20 `Deployed` and 20 `Intermediate` scale are removed according to the above rules.
- When upgrading from a release prior to 4.0(1) to release 4.0(1) or later, only the latest versions of templates are preserved.

Tagging Templates

At any point you can choose to tag the current version of the template as "golden", for example for future references to indicate a version that was reviewed, approved, and deployed with a fully validated configuration.

Procedure

- Step 1** Log in to your Nexus Dashboard Orchestrator GUI.
- Step 2** From the left navigation menu, select **Application Management > Schemas**.
- Step 3** Click the schema that contains the template you want to view.

Step 4 In the Schema view, select the template you want to review.

Step 5 From the template's actions (...) menu, select **Tag**.

If the template is already tagged, the option will change to **Un-Tag** and allows you to remove the tag from the current version.

Any version that was tagged will be indicated by a star icon in the template's version history screen.

Viewing History and Comparing Previous Versions

This section describes how to view previous versions for a template and compare them to the current version.

Procedure

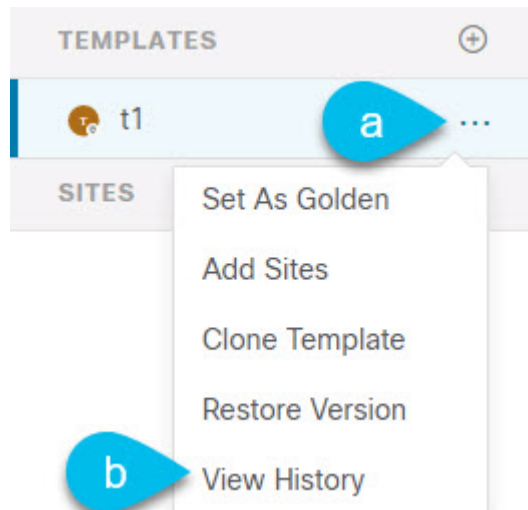
Step 1 Log in to your Nexus Dashboard Orchestrator GUI.

Step 2 From the left navigation menu, select **Application Management > Schemas**.

Step 3 Click the schema that contains the template you want to view.

Step 4 In the Schema view, select the template you want to review.

Step 5 From the template's actions (...) menu, select **History**.



Step 6 In the **Version History** window, make the appropriate selections.

Version History

General Information

- Schema versioning
- Template t1
- Tenant mso-tenant1

Versions

Golden Versions ☒ Deployed Versions ☒

Timeline: 1 — 2 — 3 — 4

Version 2 (Selected)

1 policies | 0 sites

```

1. {
2.   "siteDelta": {},
3. }
4.
5.   "anps": [],
6.   "bds": [],
7. }
34. }
```

Version 4 (Current)

2 policies | 1 sites

```

1. {
2.   "siteDelta": {
3.     "o": {
4.       "anps": [],
5.       "bds": [],
6.       "contracts": [],
7.       "externalEggs": [],
8.       "intersite3outs": [],
9.       "networks": [],
10.      "serviceGraphs": [],
11.      "siteId": "61042fa61a7b8c0a62a1a0c4",
12.      "templateName": "t1",
13.      "vrfs": []
14.    }
15.  },
16.  "template": {
17.    "anps": [],
18.    "bds": [
19.      {
20.        "arpFlood": true,
21.        "bdRef": "/schemas/610450571f0000a5030540af/templates/t1/bds/bd1",
22.        "dhcpLabels": [],
23.        "displayName": "bd1",
24.        "intersiteBumTrafficAllow": true,
25.        "l2Stretch": true

```

OK

- a) Enable the **Golden Versions** checkbox to filter the list of previous versions to display only the versions of this template that had been marked as **Golden**.

Tagging a template as "Golden" is described in [Tagging Templates, on page 56](#).

- b) Enable the **Deployed Versions** checkbox to filter the list of previous versions to display only the versions of this template that had been deployed to sites.

A new template version is created every time the template is changed and the schema is saved. You can choose to only show the versions of the template that were actually deployed to sites at some point.

- c) Click on a specific version to compare it to the current version.

The version you select is always compared to the current version of the template. Even if you filter the list using the **Golden Versions** or **Deployed Versions** filters, the current version will always be displayed even if it was never deployed or tagged as golden.

- d) Mouse over the **Edit** icon to see information about who created the version and when.

Step 7 Click **OK** to close the version history window.

Reverting Template to Earlier Version

This section describes how to restore a previous version of the template. When reverting a template, the following rules apply:

- If the target version references objects that are no longer present, restore operation will not be allowed.
- If the target version references sites that are no longer managed by NDO, restore operation will not be allowed.
- If the current version is deployed to one or more sites to which the target version was not deployed, restore operation will not be allowed.

You must first undeploy the current version from those sites before reverting the template.

- If the target version was deployed to one or more sites to which the current version is not deployed, restore operation is allowed.

Procedure

Step 1 Log in to your Nexus Dashboard Orchestrator GUI.

Step 2 From the left navigation menu, select **Application Management > Schemas**.

Step 3 Click the schema that contains the template you want to view.

Step 4 In the Schema view, select the template you want to review.

Step 5 From the **Actions (...)** menu, select **Rollback**.

Step 6 In the **Rollback** window, select one of the earlier versions to which you want to restore.

You can filter the list of versions using the **Golden Versions** and **Deployed Versions** checkboxes.

When you select a version, you can compare the template configuration of that version to the current version of the template.

Step 7 Click **Restore** to restore the selected version.

When you restore a previous version, a new version of the template is created with the same configuration as the version you selected in the previous step.

For example, if the latest template version is 3 and you restore version 2, then version 4 is created that is identical to the version 2 configuration. You can verify the restore by browsing to the template version history and comparing the current latest version to the version you had selected during restore, which should be identical.

If template review and approval (change control) is disabled and your account has the correct privileges to deploy templates, you can deploy the version to which you reverted.

However, if change control is enabled, then:

- If you revert to a version that had been previously deployed and your account has the correct privileges to deploy templates, you can immediately deploy the template.

- If you revert to a version that had not been previously deployed or your account does not have the correct privileges to deploy templates, you will need to request template approval before the reverted version can be deployed.

Additional information about review and approval process is available in the [Template Review and Approval](#), on [page 60](#) sections.

Template Review and Approval

Template review and approval (change control) workflow which allows you to set up designated roles for template designers, reviewers and approvers, and template deployers to ensure that the configuration deployments go through a validation process.

From within the NDO UI, a template designer can request review on the template they create. Then reviewers can view the history of all configuration changes for the template along with information about who made the changes and when, at which point they can approve or deny the current version of the template. If the template configuration is denied, the template designer can make any required changes and re-request review; if the template is approved, it can be deployed to the sites by a user with `Deployer` role. Finally, the deployer themselves can deny deployment of an approved template and restart the review process from the beginning.

The workflow is done at the template level, not schema level, which allows you to configure, review, and approve each template individually.

Enabling Template Approval Requirement

Before you can use the review and approval workflow for template configuration and deployment, you must enable the feature in the Nexus Dashboard Orchestrator's system settings.

Procedure

-
- Step 1** Log in to your Nexus Dashboard Orchestrator GUI.
 - Step 2** From the left navigation menu, select **Infrastructure > System Configuration**.
 - Step 3** On the **Change Control** tile, click the **Edit** icon.
 - Step 4** In the **Change Control** window, select **Enabled** to enable the feature.
 - Step 5** In the **Approvers** field, enter the number of unique approvals required before the templates can be deployed.
 - Step 6** Click **Save** to save the changes.
-

Create Users with Required Roles

Before you can use the review and approval workflow for template configuration and deployment, you must create the users with the necessary privileges in the Nexus Dashboard where the NDO service is deployed.

Procedure

Step 1 Log in to your Nexus Dashboard GUI.

Users cannot be created or edited in the NDO GUI, you must log in directly to the Nexus Dashboard cluster where the service is deployed.

Step 2 From the left navigation menu, select **Admin Console > Administrative > Users**.

Step 3 Create the required users.

The workflow depends on three distinct user roles: template designer, approver, and deployer. You can assign each role to a different user or combine the roles for the same user; users with `admin` privileges can perform all 3 actions.

There is no `Designer` role predefined on Nexus Dashboard, so the designer duties are assigned to any `Tenant Manager` or `Site Manager` user with write privileges, in addition to the default `Admin` user role:

- `Tenant Manager` should be used when the designer needs to make changes to templates associated only to a specific tenant (or a subset of tenants). In this case, the user should be mapped to the specific tenants.
- `Site Manager` should be used when the designer needs to make changes to templates that belong to different tenants.

In contrast to `Designer` role, there are pre-defined `Approver` and `Deployer` roles on the Nexus Dashboard that can be associated to the users. `Approver` and `Deployer` roles are not bound to specific tenant(s) by design. However, when creating a user role with both designer and approver (or designer and deployer) rights, follow the same guidelines as listed above.

Detailed information about configuring users and their privileges for local or remote Nexus Dashboard users is described in the [Nexus Dashboard User Guide](#).

You must have at least as many unique users with `Approver` role as the minimum number of approvals required, which you configured in [Enabling Template Approval Requirement, on page 60](#).

Note

If you disable the **Change Control Workflow** feature, any `Approver` and `Deployer` users will have read-only access to the Nexus Dashboard Orchestrator.

Requesting Template Review and Approval

This section describes how to request template review and approval.

Before you begin

You must have:

- Enabled the global settings for approval requirement, as described in [Enabling Template Approval Requirement, on page 60](#).
- Created or updated users in Nexus Dashboard with `approver` and `deployer` roles, as described in [Create Users with Required Roles, on page 60](#).
- Created a template with one or more policy configurations and assigned it to one or more sites.

Procedure

-
- Step 1** Log in to your Nexus Dashboard Orchestrator GUI as a user with `Tenant Manager`, `Site Manager`, or `Administrator` role.
- Step 2** If you assigned the `Tenant Manager` role, associate the user with the tenants.
If you used `Site Manager` or `Administrator` roles, skip this step.
If you assign the `Tenant Manager` role, you must also associate the user to the specific tenants they will manage.
- From the left navigation menu, select **Application Management > Tenants**.
 - Select the tenant which the user will manage.
 - Check the box next to the designer user you created in Nexus Dashboard.
 - Repeat this step for all other tenants the user will manage.
- Step 3** From the left navigation menu, select **Application Management > Schemas**.
- Step 4** Click the schema that contains the template for which you want to request approval.
- Step 5** In the schema view, select the template.
- Step 6** In the main pane, click **Send for Approval**.
- Note that the **Send for Approval** button will not be available in the following cases:
- The global change control option is not enabled
 - The template has no policy configurations or is not assigned to any sites
 - Your user does not have the right permissions to edit templates
 - The template has already been sent for approval
 - The template was denied by the approver user
-

Reviewing and Approving Templates

This section describes how to request template review and approval.

Before you begin

You must have:

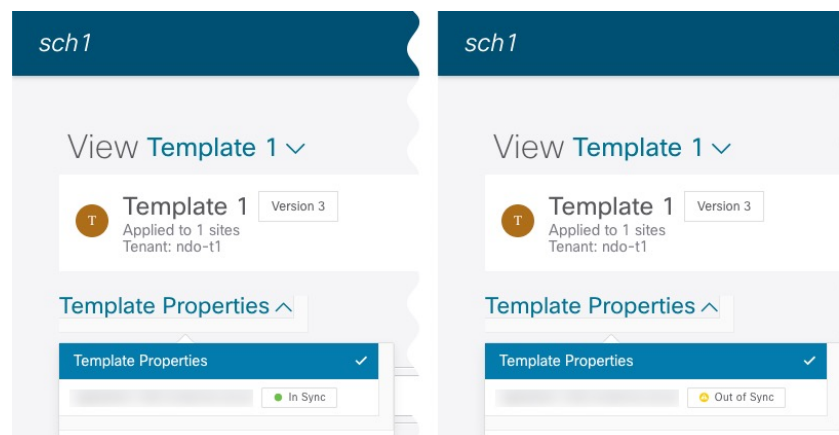
- Enabled the global settings for approval requirement, as described in [Enabling Template Approval Requirement, on page 60](#).
- Created or updated users in Nexus Dashboard with `approver` and `deployer` roles, as described in [Create Users with Required Roles, on page 60](#).
- Created a template with one or more policy configurations and assigned it to one or more sites.
- Had the template approval requested by a schema editor, as described in [Requesting Template Review and Approval, on page 61](#).

Procedure

- Step 1** Log in to your Nexus Dashboard Orchestrator GUI as a user with `Approver` or `admin` role.
- Step 2** From the left navigation menu, select **Application Management > Schemas**.
- Step 3** Click the schema that contains the template you want to review and approve.
- Step 4** In the schema view, select the template.
- Step 5** In the main pane, click **Approve**.
- If you have already approved or denied the template, you will not see the option until the template designer makes changes and re-sends the template for review again.
- Step 6** In the **Approving template** window, review the template and click **Approve**.
- The approval screen will display all the changes which the template would deploy to the sites.
- You can click **View Version History** to view the complete version history and incremental changes made between versions. Additional information about version history is available in [Viewing History and Comparing Previous Versions, on page 57](#).
- You can also click **Deployment Plan** to see a visualization and an XML of the configuration that would be deployed from this template. The functionality of the "Deployment Plan" view is similar to the "Deployed View" for already-deployed templates, which is described in [Viewing Currently Deployed Configuration, on page 69](#).

Configuration Drifts

Occasionally, you may run into a situation where the configuration actually deployed to an APIC domain is different from the configuration defined for that domain in the Nexus Dashboard Orchestrator. These configuration discrepancies are referred to as **Configuration Drifts** and are indicated by an `Out of Sync` warning next to the site name in the template view page as shown in the following figure:



Configuration drifts can manifest due to a number of different reasons. Specific steps required to resolve a configuration drift depends on its cause. Most common scenarios and their resolutions are outlined below:

- **Configuration is modified in NDO**—when you modify a template in NDO GUI, it will show as configuration drift until you deploy the changes to the sites.

To resolve this type of configuration drift, either deploy the template to apply the changes to the sites or revert the changes in the schema.

- **Configuration is modified directly in the site's APIC**—while the objects deployed from NDO are indicated by a warning icon and text in the site's APIC, an admin user can still make changes to them causing the configuration drift.

**Note**

Every time an object is modified on APIC, APIC sends a notification to Nexus Dashboard Orchestrator. On receiving the notification, Nexus Dashboard Orchestrator starts a 30 second timer (waiting for further notifications to arrive) and at the expiration of such timer then makes API calls to APIC to retrieve detailed information about the changes made all the objects for which it received a notification. This allows the Nexus Dashboard Orchestrator to display the drift symbol on the UI for all the templates where those objects are defined. The only exception to this behavior is when Nexus Dashboard Orchestrator deploys the configuration for all (or of a subset of) the objects defined in a specific template. In that case, for 60 seconds Nexus Dashboard Orchestrator would ignore any notification received from APIC relative to those specific objects and, as a consequence, it would not be able to display the drift symbol on the UI.

- **NDO configuration is restored from backup**—restoring configuration from a backup in NDO restores only the objects and their state as they were when the backup was created, it does not automatically re-deploy the restored configuration. As such, if there were changes made to the configuration and deployed on APIC since the backup was created, restoring the backup would create a configuration drift.
- **NDO configuration is restored from a backup created on an older release**—if the newer release added support for object properties which were not supported by the earlier release, these properties may cause configuration drift warning. Typically, this happens if the new properties were modified directly in the site's APIC GUI and the values are different from the defaults assumed by the Nexus Dashboard Orchestrator
- **NDO is upgraded from an earlier release**—this scenario is similar to the previous one where if new object properties are added in the new release, existing configuration may indicate a drift.

We recommend that you check for configuration drifts and, if necessary, run the "Reconcile Drift" workflow for templates, to have more visibility into the causes of the drift and be able to reconcile it. This recommendation applies to all the drift scenarios previously described in this section.

Reconciling Configuration Drifts

You can use the drift reconciliation workflow to compare the template's configuration as it is defined in Nexus Dashboard Orchestrator to the configuration rendered in the APIC controllers of the sites that are part of your Multi-Site domain. This provides better visibility into changes that may have been made in Nexus Dashboard Orchestrator or in APIC directly and give you an opportunity to correctly resolve those drifts.

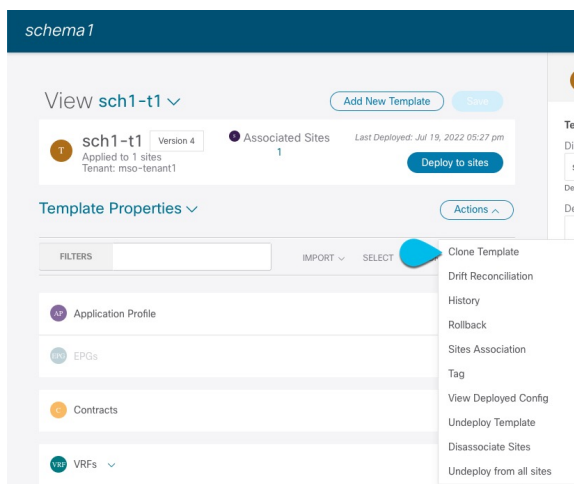
**Note**

The templates are updated and saved only after you choose **Save** or **Deploy** at the end of the reconciliation workflow. If at any time during the workflow you want to undo the changes you already chose, you can close and re-open the schema to restore the original configurations. You can then re-run the workflow from the start.

Procedure

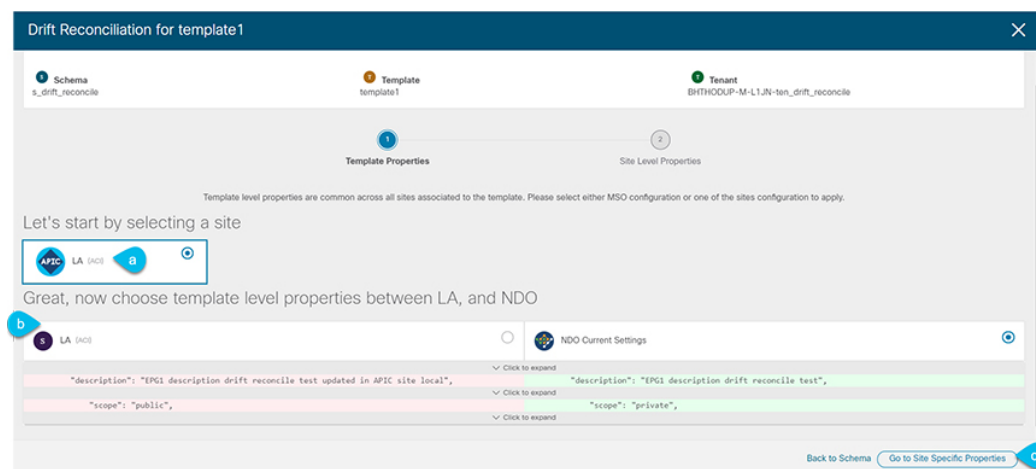
Step 1 Navigate to the schema that contains the template you want to check for configuration drifts.

Step 2 From the template's **Actions** menu, select **Reconcile Drift**.



The **Drift Reconciliation** wizard opens.

Step 3 In the **Drift Reconciliation** screen, compare the template-level configurations for each site and choose the one you want.



Template-level properties are common across all sites associated to the template. You can compare the template level properties defined on Nexus Dashboard Orchestrator with the configuration rendered in each site and decide what should become the new configuration in the Nexus Dashboard Orchestrator template. Selecting the site configuration will modify

those properties in the existing Nexus Dashboard Orchestrator template, whereas selecting the Nexus Dashboard Orchestrator configuration will keep the existing Nexus Dashboard Orchestrator template settings as is

Step 4 Click **Go to Site Specific Properties** to switch to site-level configuration.

You can choose a site to compare that specific site's configuration. Unlike template-level configurations, you can choose either the Nexus Dashboard Orchestrator-defined or actual existing configurations for each site individually to be retained as the template's site-local properties for that site.

Even though in most scenarios you will make the same choice for both template-level and site-level configuration, the drift reconciliation wizard allows you to choose the configuration defined in the site's controller at the "Template Properties" level and the configuration defined in Nexus Dashboard Orchestrator at the "Site Local Properties" level or vice versa.

Step 5 Click **Preview Changes** to verify your choices.

The preview will display full template configuration adjusted based on the choices picked in the **Drift Reconciliation** wizard. You can then click **Deploy to sites** to deploy the configuration and reconcile the drift for that template.

Cloning Schemas

This section describes how to create a copy of an existing schema and all its templates using the "Clone Schema" feature in the **Schemas** screen.



Note

If you clone a template and attempt to deploy it to the same site with a different configuration, the deployment may fail due to a duplicate name error. Changing the object name in the cloned template only updates the display name. It does not alter the database record, which causes the deployment to fail in this scenario.

Procedure

Step 1 Log in to your Nexus Dashboard Orchestrator GUI.

- Step 2** Choose the schema to clone.
- From the left navigation menu, select **Application Management > Schemas**.
 - From the **Actions** menu next to the name of the schema you want to clone, select **Clone**.
- Step 3** Provide the name for the new schema and click **Clone**.

After you click **Clone**, the UI will display `Cloning of <schema-name> was successful.` message and the new schema will be listed in the **Schemas** screen:

The new schema is created with the exact same templates (and their tenants' association), object, and policy configurations as the original schema.

Note that while the templates, objects, and configurations are copied, the site association is not preserved and you will need to re-associate the template in the cloned schema with any sites where you want to deploy them. Similarly, you will need to provide any site-specific configurations for the template objects after you associate it with the sites.

- Step 4** (Optional) Verify that the schema and all its templates were copied.
- You can verify the operation completed successfully by comparing the two schemas.

Cloning Templates

This section describes how to create a copy of an existing template using the "Clone Template" feature in the Schema view.

Procedure

- Step 1** Log in to your Nexus Dashboard Orchestrator GUI.
- Step 2** From the left navigation menu, select **Application Management > Schemas**.
- Step 3** Click the schema that contains the template you want to clone.
- Step 4** On the View menu, select a template to open it.
- Step 5** From the **Actions** menu, select **Clone Template**.
- Step 6** Provide the clone destination details.
 - From the **Destination Schema** dropdown, select the name of the Schema where you want to create the clone of the template.

You can select the same or a different schema to contain the clone of this template. If you want to clone the template into a schema that doesn't already exist, you can create a new schema by typing in the name of the schema and selecting **Create <schema-name>** option from the dropdown.

Note

When cloning across different schemas, the template must not have any objects that reference objects in other templates.

- b) In the **Cloned Template Name** field, provide the name for the new template.
- c) Click **Save** to create the clone.

A new template will be created in the destination schema, with the tenant you selected and the exact same object and policy configurations as the original template.

If the destination schema you chose was the same schema as the source template, the schema view will reload and the new template will be displayed in the left sidebar. If you chose a different schema, you can navigate to that schema to see and edit the new template.

Note that while the template objects and configurations are copied, the site association is not preserved and you will need to re-associate the cloned template with any sites where you want to deploy it. Similarly, you will need to provide any site-specific configurations for the template objects after you associate it with the sites.

Migrating Objects Between Templates

This section describes how to move objects between templates or schemas. When moving one or more objects, the following restrictions apply:

- Only EPG and Bridge Domain (BD) objects can be moved between templates.
- Migrating objects to or from Cloud Network Controller sites is not supported.
You can migrate objects between on-premises sites only.
- The source and destination templates can be in the same schema or in different schemas, but the templates must be assigned to the same tenant.
- The destination template must have been created and assigned to at least one site.
- If the destination template is not deployed and has no other objects, the template will be automatically deployed after the objects are migrated.
- Once you initiate one object migration, you cannot perform another migration that involves the same source or target template. The migration is completed when the templates have been deployed to sites.

Procedure

-
- Step 1** Log in to your Nexus Dashboard Orchestrator GUI.
 - Step 2** From the left navigation menu, select **Schemas**.
 - Step 3** Click the schema that contains the objects you want to migrate.
 - Step 4** In the Schema view, select the Template that contains the objects you want to migrate.

- Step 5** In the top right of the main pane, click **Select**.
This allows you to select one or more objects to migrate.
- Step 6** Click each object that you want to migrate.
Selected objects will display a check mark in their top right corner.
- Step 7** In the top right of the main pane, click the actions (...) icon and choose **Migrate Objects**.
- Step 8** In the **Migrate Objects** window, select the destination Schema and Template where you want to move the objects.
Only the templates with at least one site attached to them will appear in the list. If you don't see your target Template in the dropdown list, cancel the wizard and assign that template to at least one site.
- Step 9** Click **OK** and then **YES** to confirm that you want to move the objects.
The objects will be migrated from the source template to the destination template that you selected. When you deploy your configuration, the objects will be removed from any site where the source Template is deployed and added to the site where the destination template is deployed.
- Step 10** After the migration is completed, redeploy both, the source and the destination, templates.
If the destination template is not deployed and has no other objects, the template will be automatically deployed after the objects are migrated, so you can skip this step.
-

Viewing Currently Deployed Configuration

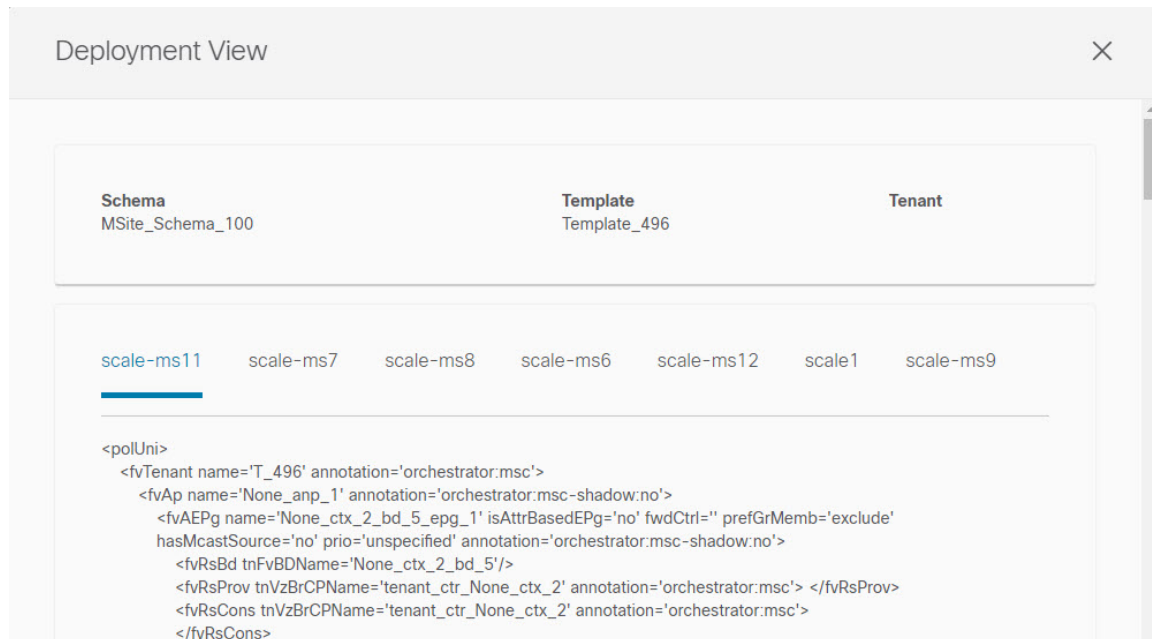
You can view all objects currently deployed to sites from a specific template. Even though any given template can be deployed, undeployed, updated, and re-deployed any number of times, this feature will show only the final state that resulted from all of those actions. For example, if `Template1` contains only `VRF1` object and is deployed to `Site1`, the API will return only `VRF1` for the template; if you then add `BD1` and redeploy, the API will return both objects, `BD1` and `VRF1`, from this point on.

This information comes from the Orchestrator database, so it does not account for any potential configuration drifts caused by changes done directly in the site's controller.

Procedure

- Step 1** Log in to your Nexus Dashboard Orchestrator GUI.
- Step 2** From the left navigation menu, select **Application Management > Schemas**.
- Step 3** Click the schema that contains the template you want to view.
- Step 4** In the left sidebar, select the template.
- Step 5** Open the **Deployed View** for the template.
a) Click the **Actions** menu next to the template's name.
b) Click **Deployed View**.
- Step 6** In the **Deployed View** screen, select the site for which you want to view the information.

You will see a graphical representation of the template configuration comparison between what's already deployed to the site and what's defined in the template..



- a) The color-coded legend indicates which objects would be created, deleted, or modified if you were to deploy the template at this time.

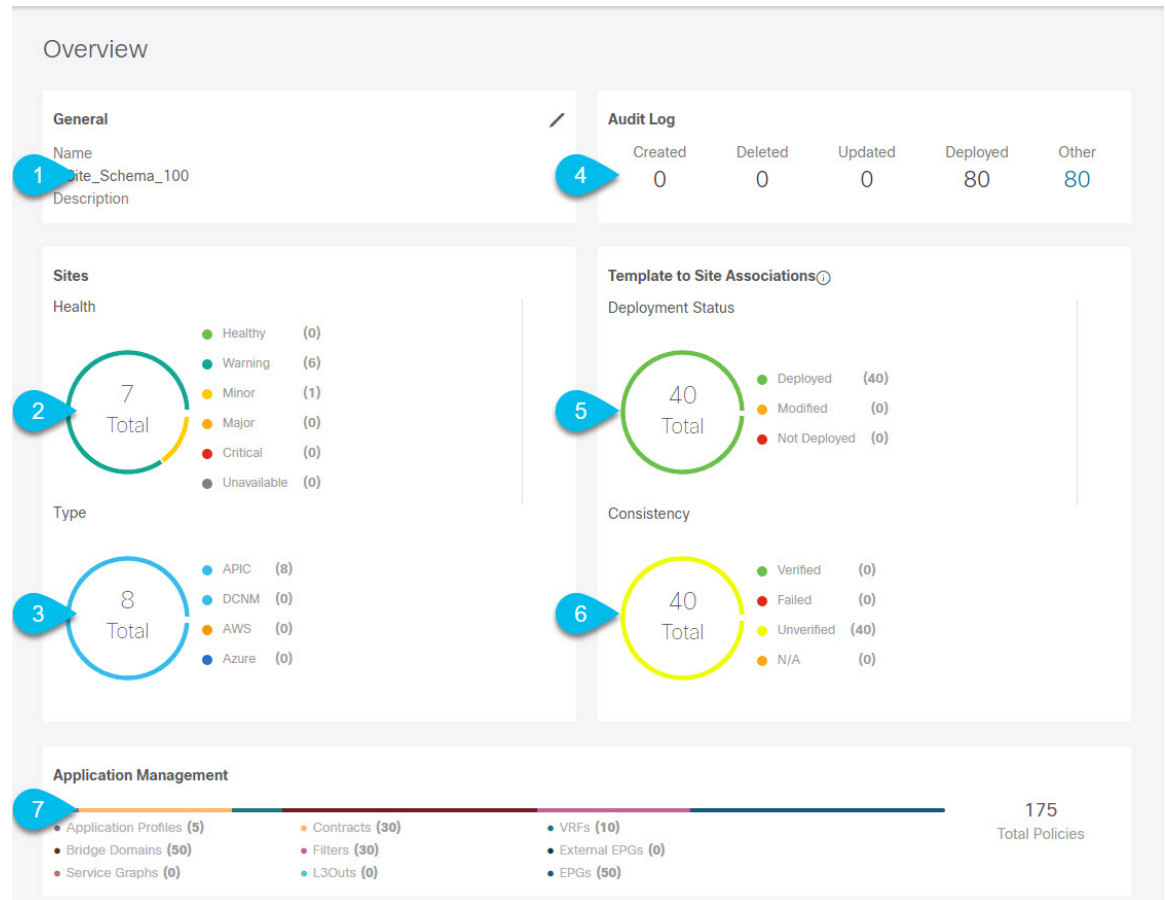
If the latest version of the template is already deployed, the view will not contain any color-coded objects and will simply display the currently deployed configuration.

- b) You can click on a site name to show configuration for that specific site.
c) You can click **View XML/JSON** to see the XML config of all the objects that are deployed to the selected site.

Schema Overview and Deployment Visualizer

When you open a schema with one or more objects defined and deployed to one or more ACI fabrics, the schema **Overview** page will provide you with a summary of the deployment.

Figure 5: Schema Overview

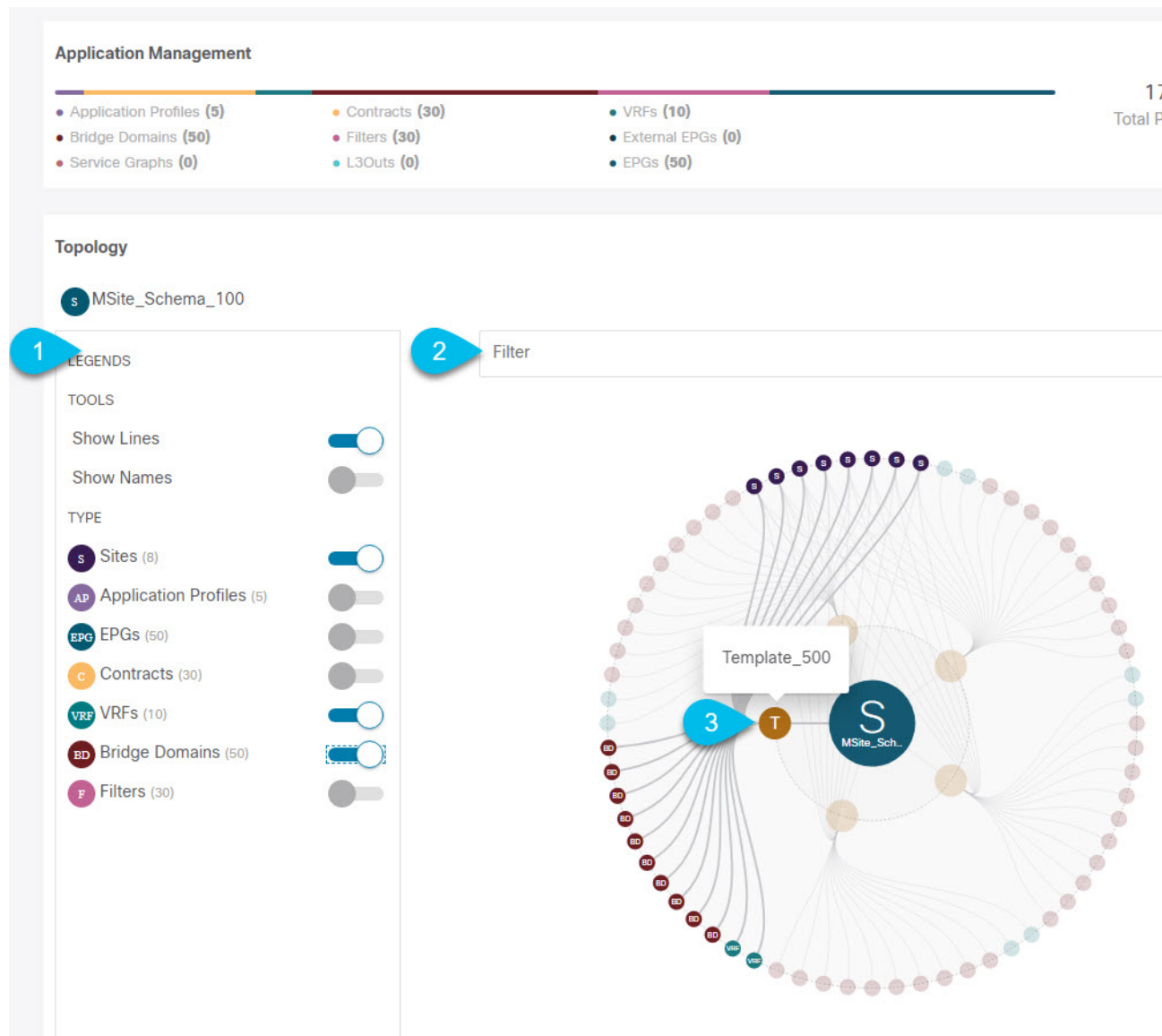


The following details are provided on this page:

- General**—Provides general information of the schema, such the name and description.
- Audit Log**—Provides audit log summary of the actions performed on the schema.
- Sites > Health**—Provides the number of sites associated with the templates in this schema sorted by the site's health status.
- Sites > Type**—Provides the number of sites associated with the templates in this schema sorted by the site's type.
- Template to Site Associations > Deployment Status**—Provides the number of templates in this schema that are associated with one or more sites and their deployment status.
- Template to Site Associations > Consistency**—Provides the number of consistency checks performed on the deployed templates and their status.
- Application Management**—Provides a summary of individual objects contained by the templates in this schema.

The **Topology** tile allows you to create a topology visualizer by selecting one or more objects to be displayed by the diagram as shown in the following figure.

Figure 6: Deployment Visualizer



1. **Legend**—Allows you to choose which policy objects to display in the topology diagram below.
2. **Filter**—Allows you to filter the displayed objects based on their names.
3. **Topology Diagram**—Provides visual representation of the policies configured in all of the Schema's templates that are assigned to sites.

You can choose which objects you want to display using the **Configuration Options** above.

You can also mouse over an objects to highlight all of its dependencies.

Finally, you can click on any object in the diagram to zoom in to see only its relationships with other objects. For example, clicking a Template will display all objects within that specific template only.



CHAPTER 4

Fabric Management

- [Creating Fabric Policies, on page 73](#)
- [Creating Fabric Resources Policies, on page 84](#)
- [Creating Monitoring Policies, on page 88](#)

Creating Fabric Policies

This section describes how to create one or more fabric policy templates. Fabric policy templates allow you to create and configure the following fabric policies:

- VLAN Pool
- Physical Domains
- SyncE Interface Policies
- Interface Settings
- Node Settings
- Pod Settings
- MACsec
- NTP Policies
- PTP Policies
- QoS DSCP Policies
- QoS SR-MPLS Policies
- QoS Class Policies

When creating Fabric Policy templates policies, consider the following:

- Fabric Policy templates do not need to be associated to any tenant, but must be map to at least one site in order to be deployed.
- The configuration of those policies is only possible at the template level and not at the specific site level.

- Undeploying a Fabric Policy template would result in preservation of the associated policies on APIC. In other words, the configuration of those policies on APIC won't be reverted to the default values or to the values that were configured on the APIC before the Orchestrator began managing them.

Procedure

Step 1 Log in to your Nexus Dashboard and open the Nexus Dashboard Orchestrator service.

Step 2 Create a new Fabric Policy template.

- From the left navigation pane, choose **Fabric Management > Fabric Policies**.
- On the **Fabric Policy Templates** page, click **Add Fabric Policy Template**.
- In the **Fabric Policies** page's right properties sidebar, provide the **Name** for the template.

By default, the new template is empty, so you need to add one or more fabric policies as described in the following steps. Note that you don't have to create every policy available in the template – you can define one or more policies of each type to deploy along with this template. If you don't want to create a specific policy, simply skip the step that describes it.

Step 3 Assign the template to one or more sites.

The process for assigning Tenant Policy templates to sites is identical to how you assign application templates to sites.

- In the **Template Properties** view, click **Actions** and choose **Sites Association**.

The **Associate Sites to <template-name>** window opens.

- In the **Associate Sites** window, check the checkbox next to the sites where you want to deploy the template.

Note that only the on-premises ACI sites support tenant policy templates and will be available for assignment.

- Click **Ok** to save.

Step 4 Create a VLAN Pool.

A VLAN pool specifies the VLAN IDs or ranges used for VLAN encapsulation that the physical or VMM domains consume.

- From the **+Create Object** dropdown, select **VLAN Pool**.
- In the right properties sidebar, provide the **Name** for the policy.
- (Optional) Click **Add Description** and provide a description for the policy.
- Click **+Add VLAN Range**, provide the range, and click the checkmark icon to save it.
- Repeat the previous substep to create any additional VLAN ranges within the same policy.
- Repeat this step to create any additional VLAN pools.

Step 5 Create a Physical Domain.

Physical domain profiles are typically used for bare metal server attachment and management access. A domain is configured to be associated with a VLAN pool. EPGs are then configured to use the VLANs associated with a domain

- From the **+Create Object** dropdown, select **Physical Domain**.
- In the right properties sidebar, provide the **Name** for the domain.
- (Optional) Click **Add Description** and provide a description for the policy.
- Click **Select a VLAN Pool Policy** and choose one of the VLAN pools for this domain.

The VLAN pool must be already created as described in Step 3.

- e) Repeat this step to create any additional Physical Domains.

Step 6

Create a SyncE Interface Policy.

With Ethernet equipment gradually replacing Synchronous Optical Networking (SONET) and Synchronous Digital Hierarchy (SDH) equipment in service-provider networks, frequency synchronization is required to provide high-quality clock synchronization over Ethernet ports. Frequency or timing synchronization is the ability to distribute precision frequency around a network. Synchronous Ethernet (SyncE) provides the required synchronization at the physical level. In SyncE, Ethernet links are synchronized by timing their bit clocks from high-quality, stratum-1-traceable clock signals in the same manner as SONET/SDH.

For detailed information on SyncE in ACI fabrics, see the "[Synchronous Ethernet \(SyncE\)](#)" chapter of the *Cisco APIC System Management Configuration Guide* for your release.

- a) From the **+Create Object** dropdown, select **SyncE Interface Policy**.
- b) In the right properties sidebar, provide the **Name** for the policy.
- c) (Optional) Click **Add Description** and provide a description for the policy.
- d) Provide policy details.
 - **Admin State** – Enables or disables the policy.
Default is `disabled`.
 - **Sync State Msg** – If unchecked, disables sending ESMC packets and also ignores any received ESMC packets.
 - **Selection Input** – Enables configuration of the priority of the frequency source on an interface.
 - **Src Priority** – The priority of the frequency source on an interface. This value is used in the clock-selection algorithm to choose between two sources that have the same QL.
Values can range from 1 (highest priority) to 254 (lowest priority). The default value is 100.
Can be configured only if **Selection Input** is enabled.
 - **Wait To Restore** – The wait-to-restore time, in minutes, is the amount of time after the interface comes up before it is used for frequency synchronization on an interface. Values can range from 0 to 12 minutes. The default value is 5.
Can be configured only if **Selection Input** is enabled.

- e) Repeat this step to create any additional SyncE Interface policies.

Step 7

Create an Interface Settings policy.

If you want to configure SyncE or MACsec for this interface, you must have those policies already created as described in corresponding steps.

An interface settings policy allows you to define a set of common interface settings which you can later deploy to one or more ports on one or more switches for consistent configuration across them.

- a) From the **+Create Object** dropdown, select **Interface Settings**.
- b) Choose the **Type** of the interface you are configuring.
- c) In the right properties sidebar, provide the **Name** for the policy.
- d) (Optional) Click **Add Description** and provide a description for the policy.
- e) Provide policy details.
 - **Speed** – The data transfer rate for the port. This should match the destination to which the port is linked. The speed can be changed only for certain ports, and not all speeds are available on all systems. For more information, see the Hardware Installation Guide for your specific switch node.

- **Auto-Negotiation** – Enables autonegotiation for port.
- **VLAN Scope** – The Layer 2 Interface VLAN scope.

Global scope – Sets the VLAN encapsulation value to map only to a single EPG per leaf.

Port Local scope – Allows allocation of separate (port, VLAN) translation entries in both ingress and egress directions. This configuration is not valid when the EPGs belong to a single bridge domain.

- **CDP Admin State** – Enables Cisco Discovery Protocol (CDP) on the interface.
- **LLDP** – Enables Link Layer Discovery Protocol (LLDP) on the interface
- **MCP Admin State** – Enables MisCabling Protocol (MCP) on the interface.
- **Domains** – Choose one or more domains with which you want to associate this interface policy.

Specifying a domain is not mandatory, the interface policy can be created and deployed to sites without an associated domain.

- **Advanced Settings** – Click the arrow next to this section to expand.

- **SyncE** – If you want have a SyncE policy defined and want to assign it to this interface settings policy, select it from the dropdown.
- **Debounce Interval** – The port debounce time is the amount of time that an interface waits to notify the supervisor of a link going down. During this time, the interface waits to see if the link comes back up.
- **Bring Up Delay** – Specifies a time in milliseconds that the decision feedback equalizer (DFE) tuning is delayed when a port is coming up. The delay is used to help avoid CRC errors during link bringup when using some third-party adapters.

You should set the delay only as required; in most cases, you do not need to set a delay.

- **FEC** – Forwarding Error Correction (FEC) is a method of obtaining error control in data transmission over an unreliable or noisy channel in which the source (transmitter) encodes the data in a redundant way using Error Correcting Code and the destination (receiver) recognizes it and corrects the errors without needing a retransmission.
- **QinQ** – Enables mapping double-tagged VLAN traffic ingressing on a regular interface, PC, or vPC to an EPG. When this feature is enabled and double-tagged traffic enters the network for an EPG, both tags are processed individually in the fabric and restored to double-tags when egressing the ACI switch. Ingressing single-tagged and untagged traffic is dropped.
- **Reflective Relay** – Forwards all traffic to an external switch, which then applies policy and sends the traffic back to the destination or target VM on the server as needed. There is no local switching. For broadcast or multicast traffic, reflective relay provides packet replication to each VM locally on the server.

One benefit of reflective relay is that it leverages the external switch for switching features and management capabilities, freeing server resources to support the VMs. Reflective relay also allows policies that you configure on the Cisco APIC to apply to traffic between the VMs on the same server.

In the Cisco ACI, you can enable reflective relay, which allows traffic to turn back out of the same port it came in on. You can enable reflective relay on individual ports, port channels, or virtual port channels as a Layer 2 interface policy.

The default value is disabled.
- **LLDP Transmit State** – Allows Link Layer Discovery Protocol (LLDP) packets to be sent from the interface.

LLDP Receive/Transmit State flags can be configured only if LLDP is globally enabled in the Interface policy.

- **LLDP Receive State** – Allows LLDP packets to be received by the interface.
- **BPDU Filter** – Bridge Protocol Data Unit (BPDU) filter filters out any BPDUs on the port.
BPDU filter prevents both inbound and outbound BPDUs – the received BPDUs are dropped and no BPDUs are sent out.
- **BPDU Guard** – BPDU guard prevents the port from receiving BPDUs; if any BPDUs are received on the port, the port is put into `errdisable` mode.
- **LLFC Transmit State** – Allows Link Level Flow Control (LLFC) packets to be sent from the interface.
- **LLFC Receive State** – Allows LLFC packets to be received by the interface.
- **Access MACsec Policy** – If you want have an access MACsec policy defined and want to assign it to this interface settings policy, select it from the dropdown.

f) Repeat this step to create any additional Interface Settings policies.

Step 8

Create a Node Settings policy.

A node settings policy allows you to define a set of common node settings which you can later deploy to one or more switches for consistent configuration across them.

In this release, the Node Settings policy supports enabling SyncE and PTP functionalities.

- a) From the **+Create Object** dropdown, select **Node Settings**.
- b) In the right properties sidebar, provide the **Name** for the policy.
- c) (Optional) Click **Add Description** and provide a description for the policy.
- d) If you want to deploy SyncE configuration to the nodes, enable **SyncE** and provide the settings.

For detailed information about SyncE, see the "Synchronous Ethernet (SyncE)" chapter of the [Cisco APIC System Management Configuration Guide](#) for your release.

- **Admin State** – Enables or disables the policy.
- **Quality Level Option** – Specifies the accuracy of the clock. This information is transmitted across the network using SSMs carried by ESMC and is used to determine the best available source to which the devices in the system can synchronize.

e) If you want to deploy PTP configuration to the nodes, enable **PTP** and provide the settings.

For detailed information about PTP, see the "Precision Time Protocol" chapter of the [Cisco APIC System Management Configuration Guide](#) for your release.

f) Repeat this step to create any additional Node Settings policies.

Step 9

Create a Pod Settings policy.

Before you can create a Pod settings policy, you must have an NTP policy already created for it as described in the corresponding step.

If you want to configure a Pod-wide MACsec policy, you must have a MACsec policy already created as described in the corresponding steps.

A Pod settings policy allows you to define a set of common Pod settings which you can later deploy to one or more Pods in your fabric for consistent configuration across them.

- a) From the **+Create Object** dropdown, select **Pod Settings**.
- b) In the right properties sidebar, provide the **Name** for the policy.
- c) (Optional) Click **Add Description** and provide a description for the policy.
- d) Click **Select a NTP Policy** and choose the NTP policy.
- e) From the **Fabric MACsec Policy** dropdown, select the MACsec policy.
- f) Repeat this step to create any additional Pod Settings policies.

Step 10

Create a MACsec policy.

MACsec provides MAC-layer encryption over wired networks by using out-of-band methods for encryption keying. The MACsec Key Agreement (MKA) Protocol provides the required session keys and manages the required encryption keys.

For detailed information on MACsec in ACI fabrics, see the "MACsec" chapter of the [Cisco APIC System Management Configuration Guide](#) for your release.

- a) From the **+Create Object** dropdown, select **MACsec**.
- b) In the right properties sidebar, provide the **Name** for the policy.
- c) (Optional) Click **Add Description** and provide a description for the policy.
- d) Provide policy details.

- **Type** – Defines the type of the interfaces this policy will be applied to.

All of the links on a spine switch are considered to be fabric links. However, if a spine switch link is used for IPN connectivity, then this link will be treated as an access link. This means that a MACsec access policy must be used to deploy MACsec on these links

- **Admin State** – Enables or disables the policy.

- **Cipher Suite** – When selecting the cipher suite AES 128 or AES 256 without Extended Packet Numbering (XPN), you must explicitly specify the Security Association Key (SAK) expiry time. Leaving the SAK expiry time value at the default ("disabled") can cause interfaces to go out of service randomly.

- **Window Size** – A replay window is necessary to support the use of MACsec over provider networks that reorder frames. Frames within the window can be received out of order, but are not replay protected. The default window size is 64. The replay window size can be configured in the range of 0 to $2^{32}-1$ if you use the Cisco APIC GUI or CLI. If you use a XPN cipher suite, the maximum replay window size is $2^{30}-1$, and if you configure a higher window size, the window size gets restricted to $2^{30}-1$. If you change the cipher suite to a non-XPN cipher suite, then there is no restriction and the configured window size is used

- **Security Policy** – APIC MACsec supports two security modes. The MACsec `Must-Secure` only allows encrypted traffic on the link while the `Should-Secure` allows both clear and encrypted traffic on the link. For example, a port can turn on MACsec in `Must-Secure` mode before its peer has received its keychain resulting in the link going down. To address this issue the recommendation is to deploy MACsec in `Should-Secure` mode and once all the links are up then change the security mode to `Must-Secure`

Note

Before deploying MACsec in `Must-Secure` mode, the keychain must be deployed on the affected interface or the interface will go down.

- **SAK Expiry Time** – When selecting the cipher suite AES 128 or AES 256 without Extended Packet Numbering (XPN), you must explicitly specify the Security Association Key (SAK) expiry time. Leaving the SAK expiry time value at the default can cause interfaces to go out of service randomly.

- **Key Name** – Allows you to create a MACsec key. APIC is responsible for the MACsec keychain distribution to all the nodes in a Pod or to particular ports on a node.
 - Click **+Add MACsec Key**.
 - Provide the **Key Name**.
 - Provide the pre-shared key in the **PSK** field.
 - In the **Start Time** field, provide the date for the key to become valid.
 - In the **End Time** field, provide the date for the key to expire.
 - Click **Ok** to save the key.
 - Repeat the steps for any additional keys you want to provide.

e) Repeat this step to create any additional MACsec policies.

Step 11

Create an NTP Settings Policy.

Within the ACI fabric, time synchronization is a crucial capability upon which many of the monitoring, operational, and troubleshooting tasks depend. Clock synchronization is important for proper analysis of traffic flows as well as for correlating debug and fault time stamps across multiple fabric nodes.

For detailed information on NTP in ACI fabrics, see the "Provisioning Core ACI Fabric Services" chapter of the [Cisco APIC Basic Configuration Guide](#) for your release.

- a) From the **+Create Object** dropdown, select **NTP Settings**.
- b) In the right properties sidebar, provide the **Name** for the policy.
- c) (Optional) Click **Add Description** and provide a description for the policy.
- d) Provide policy details.

- Click **+Add Key** to provide the NTP client authentication key.
- **Advanced Settings** – Click the arrow next to this section to expand.
 - **Admin State** – Enables or disables the NTP policy.
 - **Server State** – Enables the ACI leaf switches to act as NTP servers to provide NTP information to downstream clients.

When enabled, the downstream clients can use the in-band/out-of-band management IP address of the leaf switch to which they are connected as the NTP server.
 - **Master Mode** – Enables the designated NTP server to provide undisciplined local clock time to downstream clients with a configured stratum number. For example, a leaf switch that is acting as the NTP server can provide undisciplined local clock time to leaf switches acting as clients. This is only applicable when the server clock is undisciplined.
 - **Stratum** – Specifies the stratum level from which NTP clients will get their time synchronized.

If the **Server State** option is enabled and clients connected to an ACI leaf switch are configured to use the switch's management IP address as NTP server, they will receive NTP information with stratum+1. The range is from 1 to 14.
 - **Authentication State** – Enables certificate-based authentication.

If you enable this option, you must provide the key using the **+Add Key** option above.

- Click **+Add Provider** to specify the NTP server information.

In the **Add Provider** window that opens, you must provide the server's **Host name / IP Address**, the name of the **Management EPG**, and the **Management EPG Type**.

Note

The management EPG with the specific type that you choose must be already configured in the APICs of the sites with which this template is associated.

If you are creating multiple providers, check the **Preferred** option for the most reliable NTP source.

- Repeat this step to create any additional NTP Settings policies.

Step 12

Create a PTP Settings policy.

The Precision Time Protocol (PTP) is a time synchronization protocol for nodes distributed across a network. With PTP, you can synchronize distributed clocks with an accuracy of less than 1 microsecond using Ethernet networks. PTP's accuracy comes from the hardware support for PTP in the ACI fabric spine and leaf switches.

For detailed information on PTP in ACI fabrics, see the "[Precision Time Protocol](#)" chapter of the *Cisco APIC System Management Configuration Guide* for your release.

- From the **+Create Object** dropdown, select **PTP Settings**.
- In the right properties sidebar, provide the **Name** for the policy.
- (Optional) Click **Add Description** and provide a description for the policy.
- Provide policy details.

- **Admin State** – Enables or disables the policy.
- **Global Priority 1** – Specifies a value that is used when advertising this clock. Priority 1 overrides the default criteria (such as clock quality and clock class) for best master clock selection.
Valid values are from 0 to 255. The default is 128. Lower values take precedence.
- **Global Priority 2** – Specifies a value that is used when advertising this clock. Priority 2 is used as a tie-breaker between two devices that are otherwise equally matched in the default criteria.
Valid values are from 0 to 255. The default is 128. Lower values take precedence.
- **Global Domain** – Specifies the PTP domain number. Although multiple PTP domains are not supported in Cisco ACI, you can still change the domain number in use. The same value is used on all leaf and spine switches.
Valid values are from 0 to 128. The default is 0.
- **Fabric Profile Template** – Specifies the PTP profile that defines the default values for the interval settings below. Profiles are used to define various parameters that are optimized for different use cases of PTP. Some of those parameters include, but not limited to, the appropriate range of PTP message intervals and the PTP transport protocols. A PTP profile is defined by many organizations/standards in different industries.
 - **AES67-2015**: AES67-2015, which is the standard for audio over Ethernet and audio over IP interoperability.
 - **Default**: IEEE 1588-2008, which is the default PTP profile for clock synchronization.
 - **SMPTE-2059-2**: SMPTE ST2059-2015, which is the standard for video over IP.
 - **Telecom-8275-1**: ITU-T G.8275.1, which is the standard recommendation for telecommunications with Full Timing Support.

Full Timing Support is the term defined by ITU to describe a telecommunication network that can provide devices with PTP G.8275.1 profile on every hop. G.8275.2, which is not supported by ACI, is for Partial Timing Support that may have devices in the path that do not support PTP.

- **Fabric Announce Interval** – Specifies the logarithm of the mean interval in seconds with base 2 for a master port to send announce messages. The range depends on the chosen profile.
- **Fabric Sync Interval** – Specifies the logarithm of the mean interval in seconds with base 2 for a master port to send synchronization messages. The range and the default depend on the chosen PTP profile.
- **Fabric Delay Interval** – Specifies the logarithm of the mean interval in seconds with base 2 for a slave port to send delay request messages. The range depends on the chosen PTP profile.
- **Fabric Announce Timeout** – Specifies the number of announce messages that the system waits before the PTP announce message is considered expired. The range and the default depend on the chosen PTP profile.
- **Advanced Settings** – Click the arrow next to this section to expand.
 1. Click **+Add Profile** to add a PTP profile. Profiles are used to define various parameters that are optimized for different use cases of PTP. Some of those parameters include, but not limited to, the appropriate range of PTP message intervals and the PTP transport protocols. A PTP profile is defined by many organizations/standards in different industries
 2. In the **Add Profile** dialog, provide the **Name**.
 3. From the **Profile Template** dropdown, select one of the available profiles.
Detailed information about the profiles is available in the [Cisco APIC System Management Configuration Guide](#).
 4. Update the default profile values as required by your specific use case.

e) Repeat this step to create any additional PTP Settings policies.

Step 13

Create a QoS DSCP policy.

This policy is part of the overarching QoS preservation across IPN use case. You can use the information in this section as a reference, but we recommend following the full set of steps described in the [QoS Preservation Across IPN, on page 261](#) chapter of the *Features and Use Cases* section of this document.

- a) From the **+Create Object** dropdown, select **QoS DSCP**.
- b) In the right properties sidebar, provide the **Name** for the policy.
- c) (Optional) Click **Add Description** and provide a description for the policy.
- d) Provide policy details.

- **Admin State** – Enables or disables the policy.
- **Advanced Settings** – Click the arrow next to this section to expand.

Choose the DSCP value for each ACI QoS level. Each dropdown contains the default list of available DSCP values. You must choose a unique DSCP value for each level.

e) Repeat this step to create any additional QoS DSCP policies.

Typically, we recommend applying this policy consistently across all sites that are part of your Multi-Site domain.

Step 14

Create a QoS SR-MPLS policy.

This policy is part of the overarching SR-MPLS use case. You can use the information in this section as a reference, but we recommend following the full set of steps described in the [Multi-Site and SR-MPLS L3Out Handoff](#), on page 293 chapter of the *Features and Use Cases* section of this document.

- a) From the **+Create Object** dropdown, select **QoS SR-MPLS**.
- b) In the right properties sidebar, provide the **Name** for the policy.
- c) (Optional) Click **Add Description** and provide a description for the policy.
- d) Click **+Add Ingress Rule** to add an ingress QoS translation rule.

These rules are applied for traffic that is ingressing the ACI border leaf switches from an MPLS network and are used to map incoming packet's experimental bits (EXP) values to ACI QoS levels, as well as to set differentiated services code point (DSCP) and/or CoS values for the original traffic. To ensure that the specified CoS value is used for traffic egressing an ACI leaf node, you must also configure CoS preservation feature as part of the "QoS Class Policies".

If a custom policy is not defined or not matched, default QoS Level (Level 3) is assigned.

1. In the **Match EXP From** and **Match EXP To** fields, specify the EXP range of the ingressing MPLS packet you want to match.
2. From the **Queuing Priority** dropdown, select the ACI QoS Level to map.

This is the QoS Level you want to assign for the traffic within ACI fabric, which ACI uses to prioritize the traffic within the fabric. The options range from Level 1 to Level 6. The default value is Level 3. If you do not make a selection in this field, the traffic will automatically be assigned a Level 3 priority.

3. From the **Set DSCP** dropdown, select the DSCP value to assign to the traffic when it will be sent out of the destination ACI leaf switch.

The DSCP value specified is set in the original traffic received from the external network, so it will be re-exposed only when the traffic is VXLAN decapsulated on the destination ACI leaf node.

If you set the value to *Unspecified*, the original DSCP value of the packet will be retained.

4. From the **Set CoS** dropdown, select the CoS value to assign to the traffic when it will be sent out of the destination ACI leaf switch.

The CoS value specified is set for traffic egressing the destination ACI leaf switch. This requires CoS preservation to be enabled.

If you set the value to *Unspecified*, the original CoS value of the packet will be retained, but only if the CoS preservation option is enabled in the fabric. For more information about CoS preservation, see [Cisco APIC and QoS](#).

5. Click the checkmark icon to save the rule.
 6. Repeat these steps for any additional ingress QoS policy rules within the same policy.
- e) Click **Add Egress Rule** to add an egress QoS translation rule.

These rules are applied on the border leaf switches for the traffic that is leaving the ACI fabric via an MPLS L3Out and are used to match the DSCP value of the packet and, if a match is found, set the MPLS EXP and CoS values based on the configured policy.

If a custom policy is not defined or not matched, the default EXP value of 0 is marked on all labels. EXP values are marked in both, default and custom policy scenarios, and are done on all MPLS labels in the packet.

Custom MPLS egress policy can override existing EPG, L3Out, and Contract QoS policies.

1. Using the **Match DSCP From** and **Match DSCP To** dropdowns, specify the DSCP range you want to match for assigning the egressing MPLS packet's priority.
2. From the **Set MPLS EXP** dropdown, select the EXP value you want to assign to the egressing MPLS packet.
3. From the **Set CoS** dropdown, select the CoS value you want to assign to the egressing MPLS packet.
4. Click the checkmark icon to save the rule.
5. Repeat these steps for any additional egress QoS policy rules.

f) Repeat this step to create any additional QoS SR-MPLS policies.

Step 15

Create a QoS Class Policies policy.

Cisco ACI provides a number of user-configurable QoS levels. Cisco APIC, Release 4.0(1) and later supports 6 user-configurable QoS levels, while earlier releases supported 3. This step describes how to configure specific settings for each of these levels using the Nexus Dashboard Orchestrator.

For detailed information about QoS functionality in ACI fabrics, see [Cisco APIC and QoS](#).

The most common use case for these policies is to enable CoS preservation for traffic coming into your ACI fabric.

- a) From the **+Create Object** dropdown, select **QoS Class Policies**.
- b) In the right properties sidebar, provide the **Name** for the policy.
- c) (Optional) Click **Add Description** and provide a description for the policy.
- d) If necessary, enable **Preserve CoS**.

When traffic enters the ACI fabric, based on the configured QoS policy, each packet can be mapped to a ACI QoS level. These QoS levels are then stored in the CoS field and DE bit of the packet's outer header while the original headers are discarded. If you want to preserve the original CoS values of the ingress packets and restore it when the packet leaves the fabric, you can use this setting to enable the 802.1p Class of Service (CoS) preservation.

- e) Click **+Add Level** to define configuration details for a specific QoS class.

The **Add QoS Level Configuration** windows opens.

- f) In the **Add QoS Level Configuration** window, choose the **QoS Level** you want to configure and provide the configuration details.

- **MTU** – The maximum transmission unit to be used for packets of this QoS class.
- **Minimum Buffer** – The minimum number of reserved buffers. The number can be between 0 and 3.
The default value is 0.
- **Congestion Algorithm** – The congestion algorithm used for this QoS Level.
- **Scheduling Algorithm** – The scheduling algorithm used for this QoS Level.
- **Bandwidth Allocated** – The percentage of total bandwidth allocated to this QoS Level. The value can be between 0 and 100.
The default value is 20.
- **PFC Admin State** – The administrative state of the Priority Flow Control policy applied to FCoE traffic.
- **Admin State** – Enables or disables the policy.
- **No Drop Cos** – The CoS level to impose no drop FCoE packet handling even in case of FCoE traffic congestion.

- **PFC Scope** – The Priority Flow Control (PFC) scope. Fabric-wide PFC for the entire fabric or IntraTor PFC for the spine switches only.

g) Repeat this step to create any additional QoS Class policies.

Step 16 Click **Save** to save the changes you've made to the template.

Step 17 Click **Deploy** to deploy the template to the associated sites.

The process for deploying tenant policy templates is identical to how you deploy application templates.

If you have previously deployed this template but made no changes to it since, the **Deploy** summary will indicate that there are no changes and you can choose to re-deploy the entire template. In this case, you can skip this step.

Otherwise, the **Deploy to sites** window will show you a summary of the configuration differences that will be deployed to sites. Note that in this case only the difference in configuration is deployed to the sites. If you want to re-deploy the entire template, you must deploy once to sync the differences and then redeploy again to push the entire configuration as described in the previous paragraph.

Creating Fabric Resources Policies

This section describes how to create one or more fabric resources templates. Fabric resources templates allow you to create and configure the following:

- Physical Interfaces
- Port Channel Interfaces
- Virtual Port Channel Interfaces
- Node Profiles
- Pod Profiles

Before you begin

- Most of the fabric resource policies require one or more fabric policies, so you must have those fabric policies already defined as described in [Creating Fabric Policies, on page 73](#).

For example, when creating an interface policy (physical, port channel, or virtual port channel), you must have an interface settings policy already created.

- The templates that contain the fabric policies that are required for the fabric resources policies must be deployed before any of the fabric resource policy templates.
- The Fabric Resource Policies templates do not need to be associated to any tenant, but must be map to at least one site in order to be deployed.
- In a typical deployment, we recommend that a separate Fabric Resource Policies template is associated to each site that is part of the Multi-Site domain.

In this case, we also recommend provisioning the configuration of the associated policies always at the global template level and not at the site level.

Procedure

Step 1 Log in to your Nexus Dashboard and open the Nexus Dashboard Orchestrator service.

Step 2 Create a new Fabric Resources Policy template.

- a) From the left navigation pane, choose **Fabric Management > Fabric Resource Policies**.
- b) On the **Fabric Resource Templates** page, click **Add Fabric Resource Template**.
- c) In the **Resource Policies** page's right properties sidebar, provide the **Name** for the template.

By default, the new template is empty, so you need to add one or more fabric policies as described in the following steps. Note that you don't have to create every policy available in the template – you can define one or more policies of each type to deploy along with this template. If you don't want to create a specific policy, simply skip the step that describes it.

Step 3 Assign the template to one or more sites.

The process for assigning Tenant Policy templates to sites is identical to how you assign application templates to sites.

- a) In the **Template Properties** view, click **Actions** and choose **Sites Association**.

The **Associate Sites to <template-name>** window opens.

- b) In the **Associate Sites** window, check the checkbox next to the sites where you want to deploy the template.

Note that only the on-premises ACI sites support tenant policy templates and will be available for assignment.

- c) Click **Ok** to save.

Step 4 Create a Physical Interfaces policy.

Before you can create a physical interfaces policy, you must have an Interface Settings (**Physical**) policy already created for it as described in [Creating Fabric Policies, on page 73](#).

- a) From the **+Create Object** dropdown, select **Physical Interfaces**.
- b) In the right properties sidebar, provide the **Name** for the policy.
- c) (Optional) Click **Add Description** and provide a description for the policy.
- d) In the **Nodes** field, provide one or more node IDs where this physical interface policy will be deployed.

The configuration of the nodes policy can be also done at the site-local view of the template. In that case, the site level configuration would override the global template level configuration. As previously mentioned, in the specific scenario where a different template is created and associated to each site that is part of the Multi-Site domain, we recommend that you configure the nodes policy only at the global template level.

For example, 101, 102, 103.

- e) In the **Interfaces** field, provide the interface names where the policy will be deployed.

For example, 1/1, 1/2-4, 1/5.

- f) Click **No selected Physical Policy** and choose the interface settings policy you created for this.

The interface settings defined in the interface settings policy will be applied to the interfaces (1/1, 1/2-4, 1/5) on the nodes (101, 102, 103) you provided in the previous substeps.

- g) Repeat this step to create any additional Physical Interfaces policies.

A different policy could be needed, for example, when a unique set of physical interfaces should be configured on each node. In that case, you would define a unique Physical Interfaces policy for each specific node.

Step 5 Create a Port Channel Interfaces policy.

Before you can create a Port Channel interfaces policy, you must have an Interface Settings (PC/VPC) policy already created for it as described in [Creating Fabric Policies, on page 73](#).

- From the **+Create Object** dropdown, select **Port Channel Interfaces**.
- In the right properties sidebar, provide the **Name** for the policy.
- (Optional) Click **Add Description** and provide a description for the policy.
- In the **Node** field, provide the node ID of the switch where this physical interface policy will be deployed.

The configuration of the nodes policy can be also done at the site-local view of the template. In that case, the site level configuration would override the global template level configuration. As previously mentioned, in the specific scenario where a different template is created and associated to each site that is part of the Multi-Site domain, we recommend that you configure the nodes policy only at the global template level.

For example, 104.

- In the **Interfaces** field, provide the interface names of the interfaces that are part of the port channel.

For example, 1/6, 1/7.

- Click **No selected PC/VPC Policy** and choose the interface settings policy you created for this.

The Port Channel settings defined in the interface settings policy will be applied to the interfaces (1/6, 1/7) on the node (104) you provided in the previous substeps.

- Repeat this step to create any additional Port Channel Interfaces policies.

A different policy could be needed, for example, when a unique set of Port Channel interfaces should be configured on each node. In that case, you would define a unique Port Channel Interfaces policy for each specific node.

Step 6 Create a Virtual Port Channel Interfaces policy.

Before you can create a Virtual Port Channel interfaces policy, you must have an Interface Settings (PC/VPC) policy already created for it as described in [Creating Fabric Policies, on page 73](#).

- From the **+Create Object** dropdown, select **Virtual Port Channel Interfaces**.
- In the right properties sidebar, provide the **Name** for the policy.
- (Optional) Click **Add Description** and provide a description for the policy.
- In the **Node 1** field, provide the node ID of the first switch that contains the interfaces that are part of the virtual port channel.

For example, 105.

- In the **Interfaces on Node 1** field, provide the interfaces on the first switch.

For example, 1/8, 1/9.

- In the **Node 2** field, provide the node ID of the second switch that contains the interfaces that are part of the virtual port channel.

The configuration of the nodes policy can be also done at the site-local view of the template. In that case, the site level configuration would override the global template level configuration. As previously mentioned, in the specific scenario where a different template is created and associated to each site that is part of the Multi-Site domain, we recommend that you configure the nodes policy only at the global template level.

For example, 106.

- g) In the **Interfaces on Node 2** field, provide the interfaces on the second switch.

For example, 1/8, 1/9.

- h) Click **No selected PC/VPC Policy** and choose the interface settings policy you created for this.

The Port Channel settings defined in the interface settings policy will be applied to the interfaces on the nodes you provided in the previous substeps.

- i) Repeat this step to create any additional Virtual Port Channel Interfaces policies.

Step 7

Create a Node Profiles policy.

Before you can create a node profile policy, you must have a node settings policy already created for it as described in [Creating Fabric Policies, on page 73](#).

In this release, a node settings policy can be used to enable SyncE and/or PTP functionalities.

- a) From the **+Create Object** dropdown, select **Node Profile**.
- b) In the right properties sidebar, provide the **Name** for the policy.
- c) (Optional) Click **Add Description** and provide a description for the policy.
- d) In the **Nodes** field, provide the node IDs of the switches where you want to deploy this node profile policy.

The configuration of the nodes policy can be also done at the site-local view of the template. In that case, the site level configuration would override the global template level configuration. As previously mentioned, in the specific scenario where a different template is created and associated to each site that is part of the Multi-Site domain, we recommend that you configure the nodes policy only at the global template level.

- e) Click **No selected Node Policy** and choose the node settings policy you created for this.

The node settings defined in the node settings policy will be applied to all nodes you provided in the previous substep.

Only a single node settings policy can be referenced in a given node profile. This means that if you want to enable both SyncE and PTP policies for a given node (or set of nodes), the corresponding node settings policy with both functionalities concurrently enabled must be created (as part of a Fabric Policies template) and referenced in the node profile.

- f) Repeat this step to create any additional Node Profile policies.

Only a single Node Profile policy can be associated to a given node (or set of nodes).

Step 8

Create a Pod Profiles policy.

Before you can create a Pod profile policy, you must have a Pod settings policy already created for it as described in [Creating Fabric Policies, on page 73](#). In this release, a Pod settings policy can be used to enable the NTP functionality.

- a) From the **+Create Object** dropdown, select **Pod Profile**.
- b) In the right properties sidebar, provide the **Name** for the policy.
- c) (Optional) Click **Add Description** and provide a description for the policy.
- d) From the **Type** dropdown, select whether you want the policy to apply to **All** pods or a **Range** of pods.
- e) If you chose **Range** for the **Type**, provide the range of Pods to which to apply this policy.
- f) Click **No selected Pod Policy** and choose the Pod settings policy you created for this.

The Pod settings defined in the Pod settings policy will be applied to all nodes you provided in the previous substep.

- g) Repeat this step to create any additional Pod Profile policies.

Only a single Pod Profile policy can be associated to a given Pod (or set of Pods).

Step 9 Click **Save** to save the changes you've made to the template.

Note

When you save (or deploy) the template to one or more sites, the Orchestrator will verify that the specified nodes and/or interfaces are valid for the site(s) and will return an error if that is not the case.

Step 10 Click **Deploy** to deploy the template to the associated sites.

The process for deploying tenant policy templates is identical to how you deploy application templates.

If you have previously deployed this template but made no changes to it since, the **Deploy** summary will indicate that there are no changes and you can choose to re-deploy the entire template. In this case, you can skip this step.

Otherwise, the **Deploy to sites** window will show you a summary of the configuration differences that will be deployed to sites. Note that in this case only the difference in configuration is deployed to the sites. If you want to re-deploy the entire template, you must deploy once to sync the differences and then redeploy again to push the entire configuration as described in the previous paragraph.

Creating Monitoring Policies

This section describes how to create one or more SPAN session policies using the Monitoring Policy templates.

Before you begin

Procedure

Step 1 Log in to your Nexus Dashboard and open the Nexus Dashboard Orchestrator service.

Step 2 Create a new Tenant Policy.

- a) From the left navigation pane, choose **Fabric Management > Monitoring Policies**.
- b) On the **Monitoring Policy Templates** page, click **Add Monitoring Policy Template**.
- c) Select the SPAN session type for this template.

You can choose one of the following:

- **Tenant** – this type of SPAN sessions are usually referred to as ERSPAN sessions and allows you to configure an EPG belonging to the specified Tenant anywhere in the fabric as the SPAN session *source* and another EPG belonging to the same or to a different tenant as the *destination*.
- **Access** – allows you to configure one of the following two scenarios:
 - Access ports, port-channels, and vPC as source and destination as physical/port-channel interface, in which case the *source* and *destination* interfaces must be on the same switch.
 - Access ports, port-channels, and vPC as source and destination as EPG, in which case it is an ERSPAN session allowing the SPAN destination to be connected anywhere in the fabric.

- d) If you chose `Tenant` as the session type, choose a **Tenant** with which to associate the monitoring policy.
- e) Choose the **Site** with which to associate the monitoring policy.
- f) In the **Monitoring Policy** page's right properties sidebar, provide the **Name** for the template.

By default, the new template is empty, so you need to add one or more fabric policies as described in the following steps.

Step 3 Assign the template to one or more sites.

The process for assigning Tenant Policy templates to sites is identical to how you assign application templates to sites.

- a) In the **Template Properties** view, click **Actions** and choose **Sites Association**.

The **Associate Sites to <template-name>** window opens.

- b) In the **Associate Sites** window, check the checkbox next to the sites where you want to deploy the template.

Note that only the on-premises ACI sites support tenant policy templates and will be available for assignment.

- c) Click **Ok** to save.

Step 4 Create a SPAN Session policy for a `Tenant` type template.

If you picked `Access` for the template type, use the next step instead.

- a) From the **+Create Object** dropdown, select **SPAN Session**.
- b) In the right properties sidebar, provide the **Name** for the policy.
- c) (Optional) Click **Add Description** and provide a description for the policy.
- d) Enable the **Admin State** checkbox.

If admin state is set to `disabled`, no data is sent to the configured monitor.

- e) Click **+Add Source** and provide the SPAN source information.

For the source information, provide the following:

- **Name**
- **Direction** – the SPAN source packet direction, which can be one of the following:
 - **Both** – Replicate and forward packets that are incoming to the source and outgoing from the source.
 - **Incoming** – Replicate and forward packets that are incoming to the source.
 - **Outgoing** – Replicate and forward packets that are outgoing from the source.
- **Source EPG** – The source of the SPAN traffic.

For `Tenant` type templates, the source is always an EPG.

Click **OK** to save the source. You can then click **+Add Source** to provide additional sources if necessary.

- f) From the **Destination Group** section, provide the **Tenant**, **Destination EPG**, and **Destination IP Address** to which the replicated packets will be forwarded.

IPv4 and IPv6 IP addresses are supported in this field. However, you must not mix IPv4 for the **Destination IP** and IPv6 for the **Source IP Prefix** or vice versa.

- g) Provide the **Source IP Prefix**.

If a specific IP address is configured, all ERSPAN traffic will be sourced from that IP (for example, for all the ACI leaf switches sourcing the ERSPAN traffic). If instead a prefix is configured, each ACI leaf switch will be assigned

a unique IP that is part of that prefix to source ERSPAN traffic. This could be useful on the destination switch to distinguish the origin of the ERSPAN traffic.

h) Choose the **SPAN Version**.

i) (Optional) If necessary, configure the **Advanced Settings**.

- **Enforce SPAN Version** – When enabled, enforces the chosen SPAN Version.

If `Enabled`, the SPAN session will use the specified SPAN version if the hardware supports it. Otherwise, the session will fail.

If `Disabled` and Version 2 is specified but is not supported by the hardware, then Version 1 is used.

- **Flow ID** – The identifier of the ERSPAN packet.

When packets are copied and sent via ERSPAN, the packets are encapsulated with ERSPAN header. The flow ID is the number in the ERSPAN header to identify by which ERSPAN session these packets were copied.

The range is from 1 to 1023. The default is 1.

- **TTL** – The Time to Live (TTL) or hop limit in the 1-255 hops range; if set to zero, then no TTL is specified. The default is 64 hops.

- **DSCP** – DSCP value set in the IP header of the ERSPAN packet.

- **MTU** – The MTU of the ERSPAN-generated packets.

The range is from 64 to 9216. The default is 1518.

For ERSPAN, the real MTU received by the destination device will be larger than the configured MTU because the ERSPAN encapsulation is added. For ERSPAN version 2, an additional 46 bytes are added. For ERSPAN version 1, an additional 34 bytes are added. As a result, with the default MTU of 1518, the end device would actually need to support 1564 (1518 + 36) for version 2 or 1552 (1518 + 34) for version 1.

If the captured frame is larger than the configured MTU, then the frame is truncated to the MTU length when the frame is replicated. The packet/frame payload would be incomplete, but the headers should still be intact for analysis.

j) Repeat this step to create any additional Tenant SPAN Session policies.

Step 5

Create a SPAN Session policy for an `Access` type template.

If you picked `Tenant` for the template type, use the previous step instead.

- From the **+Create Object** dropdown, select **SPAN Session**.
- In the right properties sidebar, provide the **Name** for the policy.
- (Optional) Click **Add Description** and provide a description for the policy.
- Enable the **Admin State** checkbox.

If admin state is set to `disabled`, no data is sent to the configured monitor.

e) Click **+Add Source** and provide the SPAN source information.

For the source information, provide the following:

- **Name**

- Click **+Add Access Path** to add one or more paths on the leaf switch. The following paths are supported:

- `Port`

- Port Channel
- Virtual Port Channel
- VPC Component PC

You can use the `VPC Component PC` option if you want to configure a vPC as a source and a physical/port-channel interfaces as destination. Since for this use case all the interfaces must be on the same switch, you must not choose the vPC as a source and must select the `VPC Component PC` option representing the interfaces of that vPC on the same switch where the destination is connected. In other words, you need to create a second SPAN session for the second switch that is part of the vPC domain so that the traffic can be spanned for the source interfaces which are part of the vPC on that switch towards a local destination.

- **Direction** – the SPAN source packet direction, which can be one of the following:
 - `Both` – Replicate and forward packets that are incoming to the source and outgoing from the source.
 - `Incoming` – Replicate and forward packets that are incoming to the source.
 - `Outgoing` – Replicate and forward packets that are outgoing from the source.

- Click **+Add Filter** to provide SPAN traffic filtering information.

Traffic filtering is optional and if no filters are specified then all traffic will be spanned.

You can enable filtering based on the following attributes:

- **Src IP Prefix**
- **Src Port From**
- **Src Port To**
- **Dst IP Prefix**
- **Dst Port From**
- **Dst Port To**
- **IP Protocol**

- **SPAN Drop Packets** – Allows SPAN to capture some of dropped packets that are not captured by the regular SPAN, but it's limited to packets dropped as "forward drops".

If `Enabled`, allows for spanning of **only** dropped packets and not any traffic that was not dropped.

If `Disabled`, SPAN captures **only** the traffic that was not dropped.

The default value is `Disabled`.

- **EPG Filter** – If **SPAN Drop Packets** is disabled, you can filter the source packets based on the EPG they are coming from. To enable the filter, choose **EPG** and then select the specific EPG from the **Source EPG** dropdown.

The traffic sent and received on the previously configured source interfaces will be spanned only if it belongs to the specified EPG.

Click **OK** to save the source. You can then click **+Add Source** to provide additional sources if necessary.

- Choose the **Destination Type**.

The replicated packets can be forwarded to either an EPG or a specific access interface. In the first case, an ERSPAN session is created to send the spanned traffic to a destination connected anywhere in the fabric; in the second case, the destination must be connected to physical/port-channel interfaces on the same switch as the source interfaces.

g) If you chose `EPG` for the **Destination Type**, provide the following information:

- **Tenant, Destination EPG, and Destination IP Address** to which the replicated packets will be forwarded.

IPv4 or IPv6 IP addresses are supported in this field. However, you must not mix IPv4 for the **Destination IP** and IPv6 for the **Source IP Prefix** or vice versa.

- **Source IP Prefix** – the base IP address of the IP subnet of the source packets.

- **SPAN Version**

- (Optional) **Advanced Settings**

- **Enforce SPAN Version** – When enabled, enforces the chosen SPAN Version.

If `Enabled`, the SPAN session will use the specified SPAN version if the hardware supports it. Otherwise, the session will fail.

If `Disabled` and Version 2 is specified but is not supported by the hardware, then Version 1 is used.

- **Flow ID** – The identifier of the ERSPAN packet.

When packets are copied and sent via ERSPAN, the packets are encapsulated with ERSPAN header. The flow ID is the number in the ERSPAN header to identify by which ERSPAN session these packets were copied.

The range is from 1 to 1023. The default is 1.

- **TTL** – The Time to Live (TTL) or hop limit in the 1-255 hops range; if set to zero, then no TTL is specified. The default is 64 hops.

- **DSCP** – DSCP value set in the IP header of the ERSPAN packet.

- **MTU** – The MTU of the ERSPAN-generated packets.

The range is from 64 to 9216. The default is 1518.

For ERSPAN, the real MTU received by the destination device will be larger than the configured MTU because the ERSPAN encapsulation is added. For ERSPAN version 2, an additional 46 bytes are added. For ERSPAN version 1, an additional 34 bytes are added. As a result, with the default MTU of 1518, the end device would actually need to support 1564 (1518 + 36) for version 2 or 1552 (1518 + 34) for version 1.

If the captured frame is larger than the configured MTU, then the frame is truncated to the MTU length when the frame is replicated. The packet/frame payload would be incomplete, but the headers should still be intact for analysis.

h) Otherwise, if you chose `Access Interface` for the **Destination Type**, provide the following information instead:

- **Path Type** – type of the interface, can be `Port` or `Port Channel`.

- For `Port` interfaces, select the **Node** and **Path**.

- For `Port Channel` interfaces, select the name of the port channel.

- **MTU** – The MTU of the ERSPAN-generated packets.

The range is from 64 to 9216. The default is 1518.

For ERSPAN, the real MTU received by the destination device will be larger than the configured MTU because the ERSPAN encapsulation is added. For ERSPAN version 2, an additional 46 bytes are added. For ERSPAN version 1, an additional 34 bytes are added. As a result, with the default MTU of 1518, the end device would actually need to support 1564 (1518 + 36) for version 2 or 1552 (1518 + 34) for version 1.

If the captured frame is larger than the configured MTU, then the frame is truncated to the MTU length when the frame is replicated. The packet/frame payload would be incomplete, but the headers should still be intact for analysis.

- i) Repeat this step to create any additional Access SPAN Session policies.

Step 6 Click **Save** to save the changes you've made to the template.

Step 7 Click **Deploy** to deploy the template to the associated sites.

The process for deploying tenant policy templates is identical to how you deploy application templates.

If you have previously deployed this template but made no changes to it since, the **Deploy** summary will indicate that there are no changes and you can choose to re-deploy the entire template. In this case, you can skip this step.

Otherwise, the **Deploy to sites** window will show you a summary of the configuration differences that will be deployed to sites. Note that in this case only the difference in configuration is deployed to the sites. If you want to re-deploy the entire template, you must deploy once to sync the differences and then redeploy again to push the entire configuration as described in the previous paragraph.



PART II

Operations

- [Audit Logs, on page 97](#)
- [Backup and Restore, on page 99](#)
- [Upgrading Sites, on page 111](#)
- [Tech Support, on page 123](#)



CHAPTER 5

Audit Logs

- [Audit Logs, on page 97](#)

Audit Logs

Nexus Dashboard Orchestrator system logging is automatically enabled when you first deploy the Orchestrator cluster and captures the events and faults that occur in the environment.

You can view the Nexus Dashboard Orchestrator logs directly in the GUI by selecting **Operations > Audit Logs** from the main navigation menu.

From the **Audit Logs** page, you can click the **Most Recent** field to select a specific time period for which you want to see the logs. For example, when you select the range from November 14, 2019 to November 17, 2019 and click **Apply**, the audit log details for this time period are displayed on the **Audit Logs** page.

You can also click the **Filter** icon to filter the log details using the following criteria:

- **User:** Select this option to filter the audit logs by the user type, then click **Apply** to apply the filter.
- **Type:** Select this option to filter the audit logs by the policy types (for example, *site, user, template*) and click **Apply**.
- **Action:** Select this option to filter the audit logs by an action. The available actions are Created, Updated, Deleted, Added, Removed, Associated, Disassociated, Deployed, Undeployed, Downloaded, Uploaded, Restored, Logged in, Logged Out, Login Failed. Select an action and click **Apply** to filter the log details according to the action.



CHAPTER 6

Backup and Restore

- [Configuration Backup and Restore Guidelines, on page 99](#)
- [Configuring Remote Locations for Backups, on page 101](#)
- [Creating Backups, on page 102](#)
- [Restoring Backups, on page 102](#)
- [Exporting \(Downloading\) Backups, on page 107](#)
- [Importing Backups to Remote Location, on page 108](#)
- [Backup Scheduler, on page 108](#)

Configuration Backup and Restore Guidelines

You can create backups of your Nexus Dashboard Orchestrator configuration that can facilitate in recovering from Orchestrator failures or cluster restarts. We recommend creating a backup of the configuration before every upgrade or downgrade of your Orchestrator and after every configuration change or deployment. The backups are always created on a remote server (not Nexus Dashboard cluster), which is defined in the Nexus Dashboard Orchestrator as described in the following sections.

When creating configuration backups, the following guidelines apply:

- Importing and restoring backups created from later releases is not supported.

For example, if you downgrade your Nexus Dashboard Orchestrator to an earlier release, you cannot restore a backup of the configuration created on a later release.

- Restoring configuration backups created on releases prior to Release 4.0(1) is supported only during the initial upgrade to this release.

If you want to upgrade from a release prior to release 4.0(1) to this release, see the "Upgrading NDO Service in Nexus Dashboard" chapter in the [Cisco Nexus Dashboard Orchestrator Deployment Guide](#).

- When saving a backup, the configuration is saved in the same state in which it was deployed. When restoring a backup, any policies that were deployed will show as `deployed`, while any policies that were not deployed will remain in the `undeployed` state.
- Restoring a backup action restores the database on the Nexus Dashboard Orchestrator, but it does not make any changes to the controller (such as APIC, Cloud Network Controller, or NDFC) databases on each site.

We recommend that after you restore the Orchestrator database you resolve any configuration drifts that may appear in the templates, as described in "Configuration Drifts" section of this guide, and then

re-deploy the existing templates to avoid potentially mismatching policies between the Nexus Dashboard Orchestrator and each site's controller.

- When you create a configuration backup, the files are first created on the Orchestrator's local drives, then uploaded to the remote location, and finally deleted from the local storage. If there is not enough local disk space, the backup will fail.
- If you have a backup scheduler enabled to take local backups before upgrading to Release 4.0(1) or later, it will be disabled after the upgrade.

After the upgrade, you will need to re-add any remote locations you had set up and then re-enable backup scheduler.

- Deleting a backup using the UI also deletes the backup files from the remote location.

When restoring configuration backups, the following guidelines apply:

- If there have been no policy changes between when the backup was created and when it is being restored, no additional considerations are required and you can simply restore the configuration as described in [Restoring Backups, on page 102](#).
- If any configuration changes took place between the time when the configuration backup was created and the time it is being restored, consider the following:

- Restoring a backup will not modify any objects, policies, or configurations on the sites. Any new objects or policies created and deployed since the backup will remain deployed.

We recommend that after you restore the Orchestrator database you resolve any configuration drifts that may appear in the templates, as described in "Configuration Drifts" section of this guide, and then re-deploy the existing templates to avoid potentially mismatching policies between the Nexus Dashboard Orchestrator and each site's controller.

Alternatively, you can choose to undeploy all policies first, which will avoid any potential stale objects after the configuration is restored from backup. However, this would cause a disruption in traffic or services defined by those policies.

- The steps required to restore a configuration backup are described in [Restoring Backups, on page 102](#).
- If the configuration backup you restored was saved before it was deployed to the sites, it will be restored in the `undeployed` state and you can simply deploy it to the sites as necessary.
- If the configuration backup you restored was saved when the configuration was already deployed, it will be restored in the `deployed` state, even though none of the configurations will exist in the sites yet.

In this case, resolve any configuration drifts that may appear in the templates, as described in "Configuration Drifts" section of this guide and re-deploy the templates to sync the Nexus Dashboard Orchestrator's configuration with the sites.

- If sites that were managed when the backup was created are no longer present in the Nexus Dashboard, the restore will fail.
- If sites' status since the backup has changed (`managed` vs `unmanaged`) but the sites are still present in the Nexus Dashboard, the status will be restored to what it was at the time of backup.

Configuring Remote Locations for Backups

This section describes how to configure a remote location in Nexus Dashboard Orchestrator to which you can then export your configuration backups.

Procedure

Step 1 Log in to your Nexus Dashboard and open the Nexus Dashboard Orchestrator service.

Step 2 From the left navigation pane, select **Operations > Remote Locations**.

Step 3 In the top right of the main window, click **Add Remote Location**.

An **Add New Remote Location** screen appears.

Step 4 Provide the name for the remote location and an optional description.

Two protocols are currently supported for remote export of configuration backups:

- SCP
- SFTP

Note

SCP is supported for non-Windows servers only. If your remote location is a Windows server, you must use the SFTP protocol

Step 5 Specify the host name or IP address of the remote server.

Based on your **Protocol** selection, the server you specify must allow SCP or SFTP connections.

Step 6 Provide the full path to a directory on the remote server where you will save the backups.

The path must start with a slash (/) characters and must not contain periods (.) or backslashes (\). For example, */backups/multisite*.

Note

The directory must already exist on the remote server.

Step 7 Specify the port used to connect to the remote server.

By default, port is set to 22.

Step 8 Specify the authentication type used when connecting to the remote server.

You can configure one of the following two authentication methods:

- **Password**—provide the username and password used to log in to the remote server.
- **SSH Private Files**—provide the username and the SSH Key/Passphrase pair used to log in to the remote server.

Step 9 Click **Save** to add the remote server.

Creating Backups

This section describes how to create a new backup of your Nexus Dashboard Orchestrator configuration.

Before you begin

You must first add the remote location as described in [Configuring Remote Locations for Backups](#), on page 101.

Procedure

Step 1 Log in to your Nexus Dashboard Orchestrator.

Step 2 Backup existing deployment configuration.

- a) From the left navigation pane, select **Operations > Backups & Restore**.
- b) In the main window, click **New Backup**.

A **New Backup** window opens.

- c) Provide the backup information.

- In the **Name** field, provide the name for the backup file.

The name can contain up to 10 alphanumeric characters, but no spaces or underscores (_).

- From the **Remote Location** drop-down, select a remote location you have configured for storing backups.

- (Optional) In the **Remote Path**, provide the specific directory on the remote server where to save the backup.

The directory you specify must already exist.

- d) Click **Save** to create the backup.
-

Restoring Backups

This section describes how to restore a Nexus Dashboard Orchestrator configuration to a previous state.

Before you begin

- You must have configured a remote location for storing your NDO backups, as described in [Configuring Remote Locations for Backups](#), on page 101.
- Ensure that the backup you want to restore is on the remote location server or import the backup into the remote location, as described in [Importing Backups to Remote Location](#), on page 108.



Note Restoring a backup action restores the database on the Nexus Dashboard Orchestrator, but it does not make any changes to the controller (such as APIC, Cloud Network Controller, or NDFC) databases on each site.

We recommend that after you restore the Orchestrator database you resolve any configuration drifts that may appear in the templates, as described in "Configuration Drifts" section of this guide, and then re-deploy the existing templates to avoid potentially mismatching policies between the Nexus Dashboard Orchestrator and each site's controller.

For information on specific configuration mismatch scenarios and recommended restore procedures related to each one, see [Configuration Backup and Restore Guidelines, on page 99](#).

Procedure

Step 1 Log in to your Nexus Dashboard Orchestrator GUI.

Step 2 If necessary, undeploy existing policies.

We recommend you perform this step if new objects or policies were added to the configuration between when the backup was created and current configuration. Additional context is available in [Configuration Backup and Restore Guidelines, on page 99](#).

Step 3 From the left navigation menu, select **Operations > Backups & Restore**.

Step 4 In the main window, click the actions (...) icon next to the backup you want to restore and select **Rollback to this backup**.

If the version of the selected backup is different from the running Nexus Dashboard Orchestrator version, the rollback could cause a removal of the features that are not present in the backup version.

Step 5 Click **Yes** to confirm that you want to restore the backup you selected.

If you click **Yes**, the system terminates the current session and the user is logged out.

Note

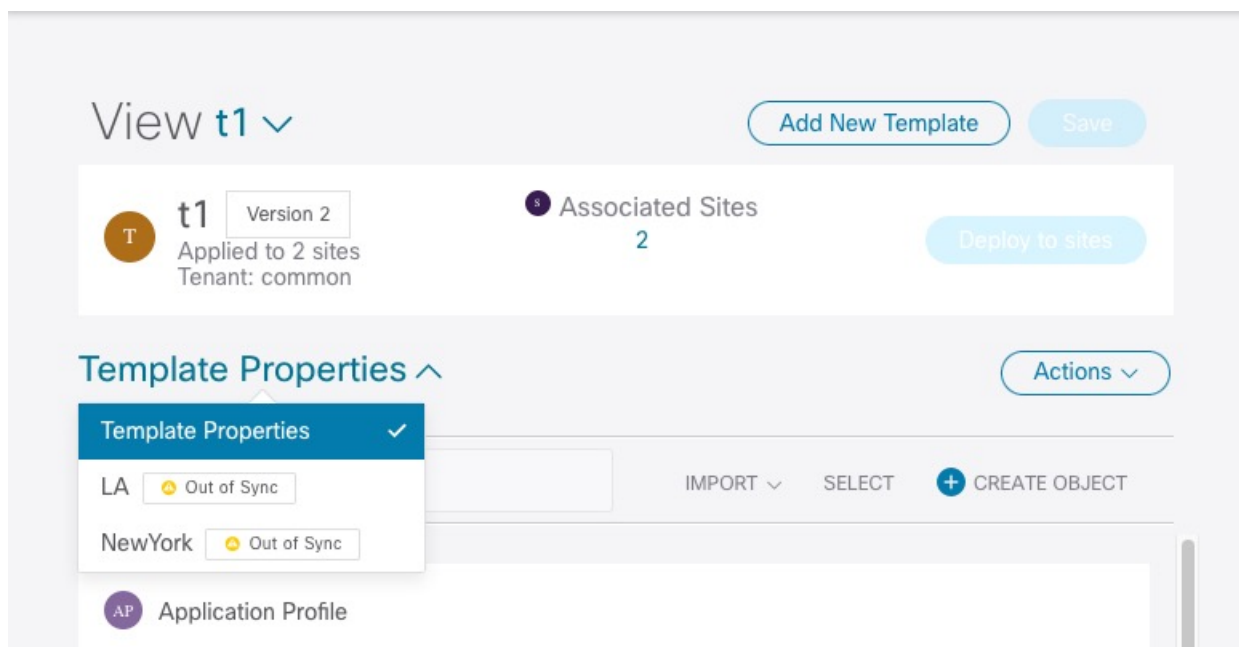
Multiple services are restarted during the configuration restore process. As a result, you may notice an up to 10 minute delay before the restored configuration is properly reflected in the NDO GUI.

Step 6 Check if any templates contain configuration drifts.

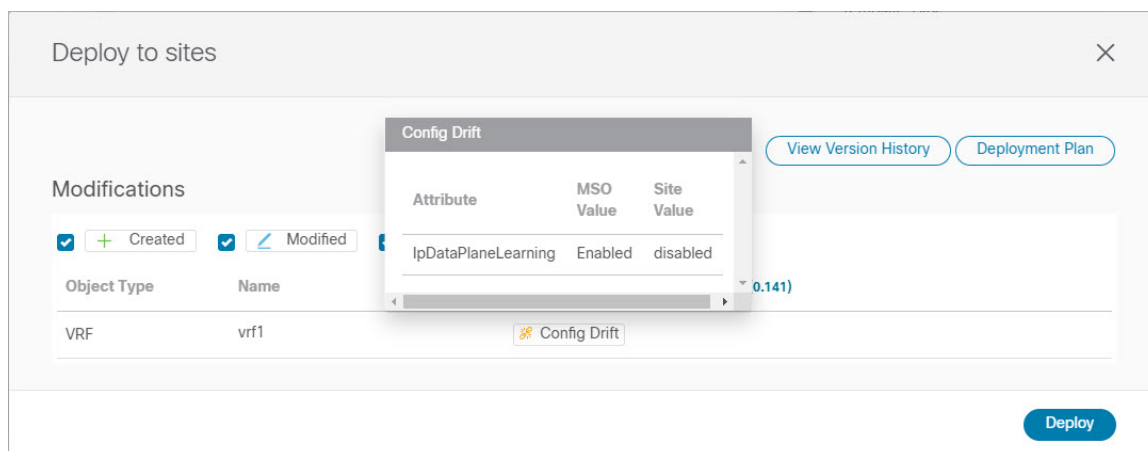
You will repeat the following steps for every schema and template in your deployment

You can check for configuration drifts in one of the following two ways:

- Check the template deployment status icon for each site to which the template is assigned:



- Select the template and click **Deploy to sites** to bring up the configuration comparison screen to check which objects contain configuration drifts:



Step 7

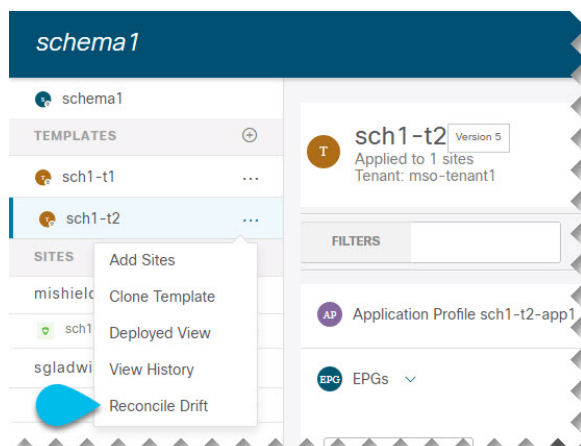
If any template contains a configuration drift, resolve the conflicts.

For more information about configuration drifts, check the "Configuration Drifts" chapter in the [Cisco Nexus Dashboard Orchestrator Configuration Guide for ACI Fabrics](#).

- Close the template deployment dialog to return to the Schema view.

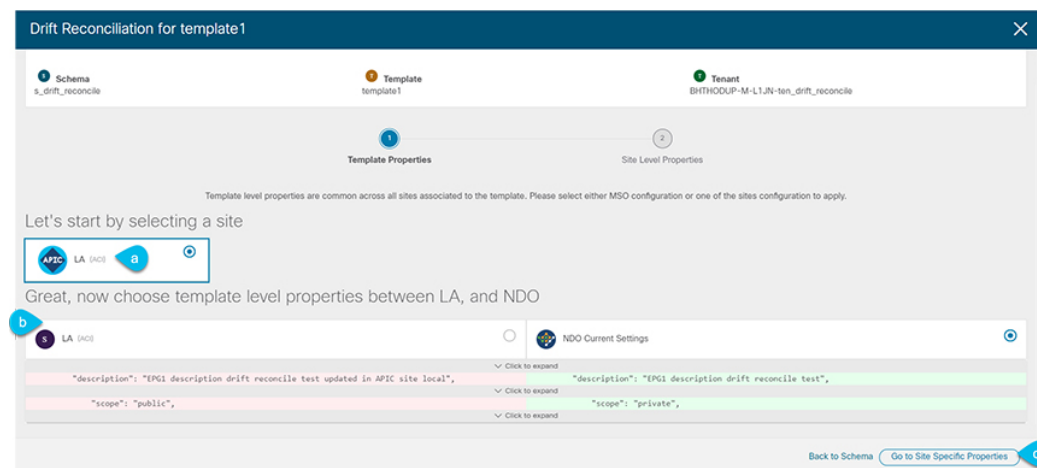
Deploying any templates at this point would push the values in the Orchestrator database and overwrite any existing settings in the fabrics.

- From the template's **Actions** menu, select **Reconcile Drift**.



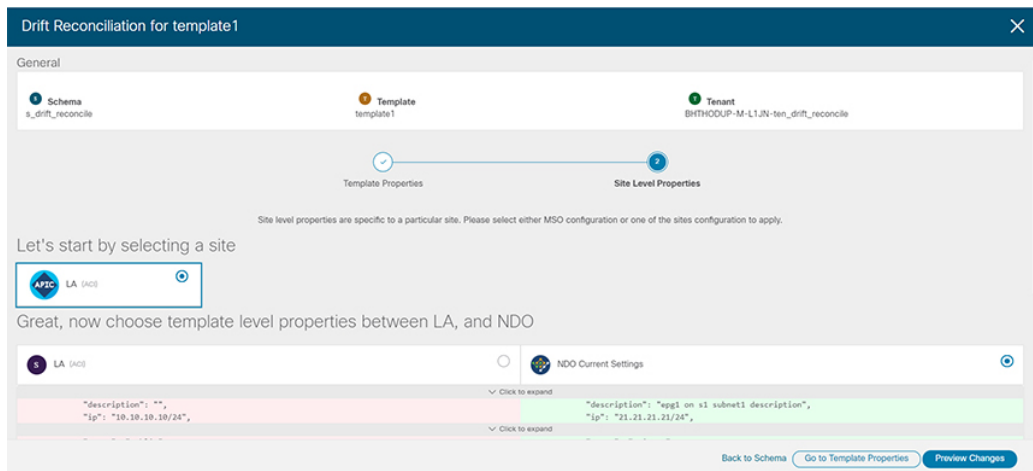
The **Drift Reconciliation** wizard opens.

- c) In the **Drift Reconciliation** screen, compare the template-level configurations for each site and choose the one you want.



Template-level properties are common across all sites associated to the template. You can compare the template level properties defined on Nexus Dashboard Orchestrator with the configuration rendered in each site and decide what should become the new configuration in the Nexus Dashboard Orchestrator template. Selecting the site configuration will modify those properties in the existing Nexus Dashboard Orchestrator template, whereas selecting the Nexus Dashboard Orchestrator configuration will keep the existing Nexus Dashboard Orchestrator template settings as is

- d) Click **Go to Site Specific Properties** to switch to site-level configuration.



You can choose a site to compare that specific site's configuration. Unlike template-level configurations, you can choose either the Nexus Dashboard Orchestrator-defined or actual existing configurations for each site individually to be retained as the template's site-local properties for that site.

Even though in most scenarios you will make the same choice for both template-level and site-level configuration, the drift reconciliation wizard allows you to choose the configuration defined in the site's controller at the "Template Properties" level and the configuration defined in Nexus Dashboard Orchestrator at the "Site Local Properties" level or vice versa.

- e) Click **Preview Changes** to verify your choices.

The preview will display full template configuration adjusted based on the choices picked in the **Drift Reconciliation** wizard. You can then click **Deploy to sites** to deploy the configuration and reconcile the drift for that template.

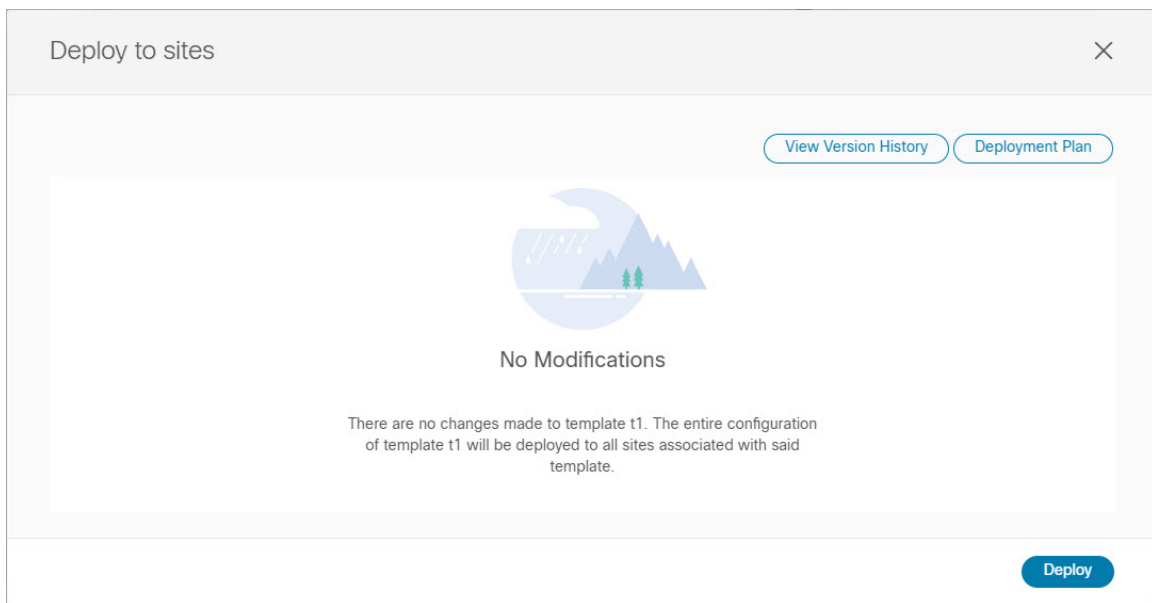
Step 8

After all configuration drifts are resolved and there are no changes shown in the **Deploy to sites** dialog for the template, perform full redeployment of the template.

Note

Due to database transformations in Release 3.7(1), you must perform a full redeployment of each template.

Ensure that the **Deploy to sites** dialog contains no changes as shown in the following figure, then click **Deploy** to redeploy complete configuration:



Step 9 Repeat the above steps for every schema and template in your Nexus Dashboard Orchestrator.

Step 10 Check audit logs to verify that all templates have been re-deployed.

You can view the audit logs in the **Operations** tab.

Audit Logs page and confirm that all templates show as `Redeployed` to ensure that full re-deployment successfully completed.

Exporting (Downloading) Backups

This section describes how to download the backup from the Nexus Dashboard Orchestrator.

Before you begin

Procedure

Step 1 Log in to your Nexus Dashboard Orchestrator GUI.

Step 2 From the left navigation menu, select **Operations > Backups & Restore**.

Step 3 In the main window, click the actions (...) icon next to the backup you want to download and select **Download**.

This will download the backup file in `msc-backups-<timestamp>.tar.gz` format to your system. You can then extract the file to view its contents.

Importing Backups to Remote Location

This section describes how to upload an existing configuration backup you have previously downloaded and import it into one of the remote locations configured in your Nexus Dashboard Orchestrator.

Before you begin

You must have completed the following:

- Created and downloaded a configuration backup as described in [Creating Backups, on page 102](#) and [Exporting \(Downloading\) Backups, on page 107](#).

If your backup is already on a remote location, for example if it was created on release 3.4(1) or later, you can download it to your local machine and upload it to a different remote location.

- Added a remote location for backups as described in [Configuring Remote Locations for Backups, on page 101](#).

Procedure

-
- Step 1** Log in to your Nexus Dashboard Orchestrator.
- Step 2** From the left navigation pane, select **Operations > Backups & Restore**.
- Step 3** In the main pane, click **Upload**.
- Step 4** In the **Upload from file** window that opens, click **Select File** and choose the backup file you want to import.
Uploading a backup will add it to the list of the backups displayed the **Backups** page.
- Step 5** From the **Remote Location** dropdown menu, select the remote location.
- Step 6** (Optional) Update the remote location path.
The target directory on the remote server, which you configured when creating the remote backup location, will be displayed in the **Remote Path** field.
You can choose to append additional subdirectories to the path. However, the directories must be under the default configured path and must have been already created on the remote server.
- Step 7** Click **Upload** to import the file.
Importing a backup will add it to the list of the backups displayed the **Backups** page.
Note that even though the backups are shown on the NDO UI, they are located on the remote servers only.
-

Backup Scheduler

This section describes how to enable or disable the backup scheduler, which will perform complete configuration backup at regular intervals.

Before you begin

You must have already added a remote location for backups as described in [Configuring Remote Locations for Backups](#), on page 101.

Procedure

- Step 1** Log in to your Nexus Dashboard Orchestrator GUI.
- Step 2** From the left navigation menu, select **Operations > Backups & Restore**.
- Step 3** In the top right of the main pane, click **Scheduler**.
The **Backup Scheduler Settings** window will open.
- Step 4** Set up backup scheduler.
- a) Check the **Enable Scheduler** checkbox.
 - b) In the **Select Starting Date** field, provide the day when you want the scheduler to start.
 - c) In the **Select Time** fields, provide the time of day when you want the scheduler to start.
 - d) From the **Select Frequency** dropdown, choose how often the backup should be performed
 - e) From the **Remote Location** dropdown, select the location where the backups will be saved.
 - f) (Optional) In the **Remote Path** field, update the path on the remote location where the backups will be saved.

The target directory on the remote server, which you configured when creating the remote backup location, will be displayed in the **Remote Path** field.

You can choose to append additional subdirectories to the path. However, the directories must be under the default configured path and must have been already created on the remote server.
 - g) Click **OK** to finish.
- Step 5** If you want to disable the backup scheduler, simply uncheck the **Enable Scheduler** checkbox in the above step.
-



CHAPTER 7

Upgrading Sites

- [Overview, on page 111](#)
- [Guidelines and Limitations, on page 113](#)
- [Downloading Controller and Switch Node Firmware to Sites, on page 113](#)
- [Upgrading Controllers, on page 116](#)
- [Upgrading Nodes, on page 118](#)

Overview



Note This feature is supported for Cisco APIC sites only. It is not supported for Cisco Cloud Network Controller or Cisco NDFC fabrics.

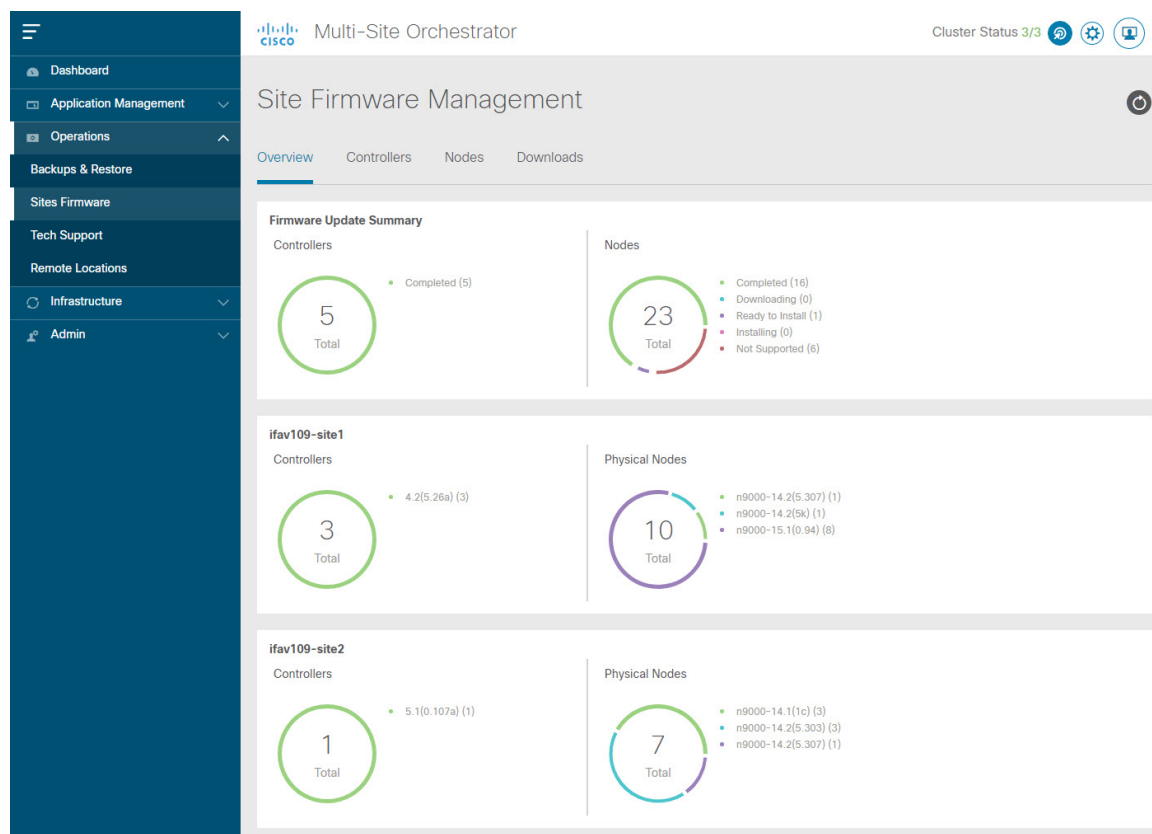
Prior to Release 3.1(1), when you deployed Cisco Multi-Site, each site's APIC clusters and switch nodes software had to be managed individually at the site level. As the number of sites in your Multi-Site domain grew, the release life cycle and upgrades could become complicated as they had to be manually coordinated and managed for release and feature compatibility.

Starting with Release 3.1(1), Cisco Nexus Dashboard Orchestrator provides a workflow that allows you to manage all sites' software upgrades from a single point eliminating the need for multiple site administrators to manually coordinate software upgrades and giving you insight into any potential issues that could affect the upgrades.

You can access the site upgrades screen by navigating to **Operations > Sites Firmware**. The page contains four tabs, which are described in this and following sections.

The **Overview** tab displays information about the sites in your Multi-Site domain and the firmware versions that are deployed or ready to be deployed. The `Sites Firmware` service polls the sites every 5 minutes for new or changed data such as the latest status of any of the upgrade policies. You can manually trigger a refresh by clicking the **Refresh** button in the upper right corner of the main pane.

Figure 7: Sites Firmware Overview



The page is divided into the following areas:

- **Firmware Update Summary**—provides overall summary of the firmware images that are present across all sites in your Multi-Site domain, including the Cisco APIC and the switch firmware.

For each type of image, the specific information includes the number of images in each state:

- **Completed**—the image is currently deployed to the controllers or the switches.
- **Downloading** (for switch nodes only)—the image is being downloaded to the switch nodes.
- **Ready to Install** (for switch nodes only)—the image was successfully downloaded to the switch nodes and is ready to be installed.
- **Installing**—the images currently in the process of being deployed to the controllers or the switch nodes.
- **Not Supported**—the images that do not support remote firmware upgrades, such as releases prior to Release 4.2(5).
- **Site-specific information**—additional sections of the page display information about individual sites, which includes the version of the currently deployed software and the number of controllers or nodes.

Guidelines and Limitations

When performing fabric upgrades from the Nexus Dashboard Orchestrator, the following restrictions apply:

- You must review and follow the guidelines, recommendations, and limitations specific to the Cisco APIC upgrade process described in the [Upgrading and Downgrading the Cisco APIC and Switch Software](#) of the *Cisco APIC Installation, Upgrade, and Downgrade Guide*.
- Your Nexus Dashboard Orchestrator must be deployed in Cisco Nexus Dashboard.

The site upgrade feature is not available for NDO deployments in VMware ESX and you need to follow the standard upgrade procedures described in [Cisco APIC Installation, Upgrade, and Downgrade Guide](#)

- The fabrics must be running Cisco APIC, Release 4.2(5) or later.

Fabrics running earlier APIC releases will not be available for selection during the upgrade workflow. Follow the standard upgrade procedures described in [Cisco APIC Installation, Upgrade, and Downgrade Guide](#).

- We recommend coordinating the site upgrades with the site administrators managing those fabrics. You may need access to the controllers or switch nodes to troubleshoot any potential issues should they arise.
- If a fabric switch node goes into an `inactive` state in the middle of the upgrade process, for example due to hardware or power failure, the process will be unable to complete. You will not be able to remove or modify the node upgrade policy from NDO during this time as NDO is unable to differentiate whether the node went down or is simply in the middle of a reboot for the upgrade.

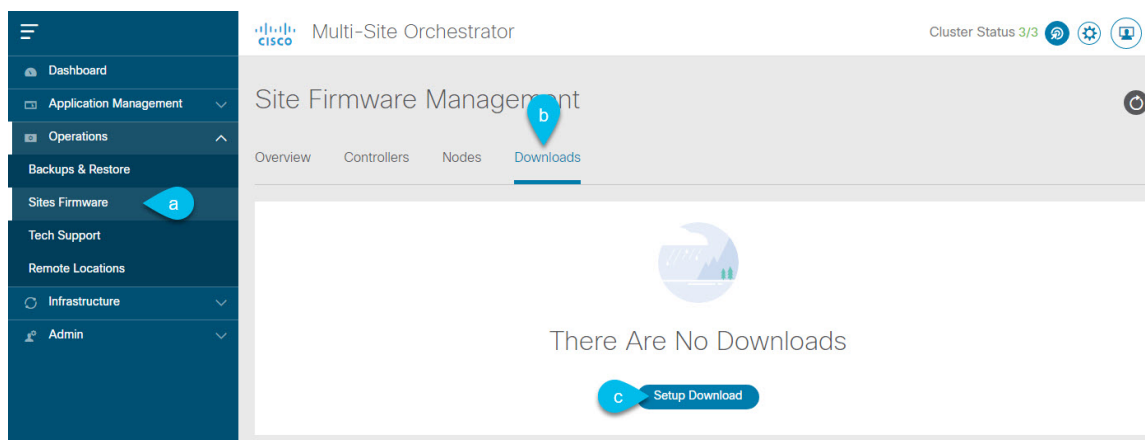
To resolve this issue, you must manually decommission the inactive node from the APIC, at which point the NDO upgrade policy will recognize the change and return a `failed` status. Then you can update the upgrade policy on the NDO to remove the switch and re-run the upgrade.

Downloading Controller and Switch Node Firmware to Sites

You must download the controller and switch software to all the site controllers in your fabrics before performing the upgrade. After you complete the following steps, you will be able to start the upgrade process at a later time using the downloaded images.

Procedure

-
- Step 1** Log in to your Nexus Dashboard Orchestrator.
 - Step 2** Set up firmware download.



- a) From the left navigation pane, select **Operations** > **Sites Firmware**.
- b) In the main window, select the **Downloads** tab.
- c) Click **Setup Downloads** tab.

If you have previously set up one or more downloads, click the **Setup Downloads** button in the top right of the main pane instead.

The **Download Image to APIC** screen opens.

Step 3 Select the sites.

The image will be downloaded to the Cisco APICs of all the sites you select here.

- a) Click **Select Sites**.
- b) In the **Select Sites** window, check one or more sites and click **Add and Close**.
- c) Click **Next** to proceed.

Step 4 Provide the download details.

- a) Provide the **Name**.

You can provide a descriptive name for tracking the download.

- b) Choose the protocol.

You can choose to download the image via **HTTP** or **SCP**.

- c) Click + **Add URL** to provide location of the image(s).

You can provide both, the APIC and the switch firmware images.

- d) If you selected **SCP**, provide the authentication information.

You need to provide the login **Username**, for example `admin`.

Then choose the **Authentication Type**:

- For **Password** authentication, simply enter the password for the username you provided earlier.
- For **SSH Key** authentication, you will need to enter the **SSH Key** and the **SSH Key Passphrase**.

- e) Click **Next** to proceed.

Step 5

In the confirmation screen, review the information and click **Submit** to proceed.

In the **Downloading** screen that opens, you can view the status of the image download.

You can also click the status, to see additional details about the progress.

Image Download - MSO-d11

Setup | **Downloading** | Complete

Download Details

Name: MSO-d11

Overall Status: **Downloading**

Status Breakdown: 3 (Downloaded: 0, Downloading: 3, Download Failed: 0)

Sites

Filter by attributes

Site	URLs	Status
ifav109-site1	1	Downloading (1)
ifav109-site2	1	Downloading (1)
ifav109-site3	1	Downloading (1)

ifav109-site3

URLs

Link: 0.117a final aci-apic-dk9.5.1.0.117a.iso

Status: Downloading (30%)

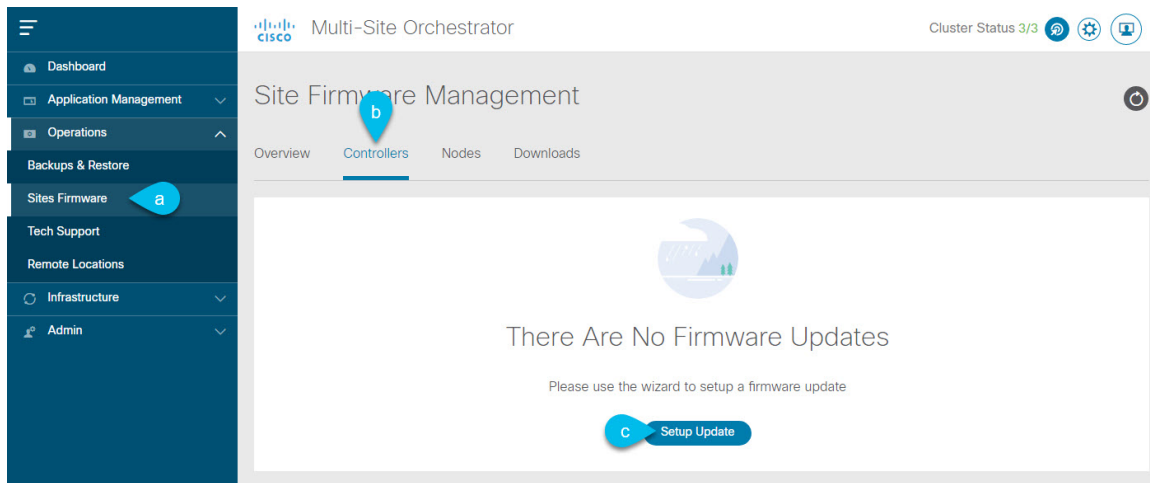
After all downloads complete, you will transition to the **Completed** screen. You do not have to wait at the **Downloading** screen, you can always navigate back to it from the **Downloads** tab by clicking the download name you provided in a previous step.

Upgrading Controllers

This section describes how to set up a software upgrade for your sites' APIC clusters.

Procedure

- Step 1** Log in to your Nexus Dashboard Orchestrator.
- Step 2** Set up APIC cluster upgrade.



- a) From the left navigation pane, select **Operations** > **Sites Firmware**.
- b) In the main window, select the **Controllers** tab.
- c) Click **Setup Update** tab.

If you have previously set up one or more updates, click the **Setup Update** button in the top right of the main pane instead.

The **Setup Site Firmware Update** screen opens.

Step 3 Provide the upgrade details.

- a) Provide the **Name**.
This is the controller upgrade policy name you will be able to use to track the upgrade progress at any time.
- b) Click **Select Sites**.
The **Select Sites** window opens.
- c) In the **Select Sites** window, check one or more sites and click **Add and Close**.
- d) Click **Next** to proceed.

Step 4 In the **Version Selection** screen, select the firmware version and click **Next**.

The firmware must be downloaded to the sites before it becomes available here. If the download you set up in previous section has completed successfully but the image is still not available here, close the **Setup Site Firmware Update** screen, navigate back to **Operations** > **Sites Firmware** > **Controllers** tab, and click the **Refresh** button to reload the latest information available for the sites; then restart the upgrade steps.

Step 5 In the **Validation** screen, review the information, then click **Next**.

Ensure that there are no faults and review any additional information that may affect your upgrade:

Setup Site Firmware Update

Setup
Downloading
Ready to Install
Installing
Complete

Site Selection
Version Selection
Validation
Confirmation

ifav109-site1

Following nodes are not in vPC ['1111','102','101','104','103'].
Configure vPC for the listed leaf nodes to avoid traffic loss during the reboot of leaf nodes.

ifav109-site1

Pod(s) [2] have fewer than two route reflectors for infra MP-BGP.
Configure spine nodes as route reflector for infra MP-BGP. Make sure that at least one route reflector spine is always up by upgrading/downgrading them in separate groups.

ifav109-site3

Following nodes are not in vPC ['301','302'].
Configure vPC for the listed leaf nodes to avoid traffic loss during the reboot of leaf nodes.

ifav109-site3

Pod(s) [1] have fewer than two route reflectors for infra MP-BGP.
Configure spine nodes as route reflector for infra MP-BGP. Make sure that at least one route reflector spine is always up by upgrading/downgrading them in separate groups.

ifav109-site3

NTP is not configured.
Configure NTP via System > QuickStart > First time setup of the ACI fabric > NTP. This is recommended to avoid any issues in database synchronization between nodes, SSL certificate check, etc.

ifav109-site3

APICs are not running recommended CIMC versions :node-1: 4.0(2f)
Upgrade to the recommended CIMC version. APICs have recommended CIMC versions based on its hardware model and APIC firmware version.

Previous
Next

Step 6 In the **Confirmation** screen, review the information and click **Submit** to start the upgrade.

Step 7 In the **Ready to Install** screen, click **Install** to start the upgrade.

If NDO to site connectivity is lost during the upgrade process, the GUI will display the last known status of the upgrade prior to loss of connectivity. Once connectivity is re-established, the upgrade status will be refreshed. You can perform a manual refresh after connectivity loss by clicking the **Refresh** button in the top right of the main pane.

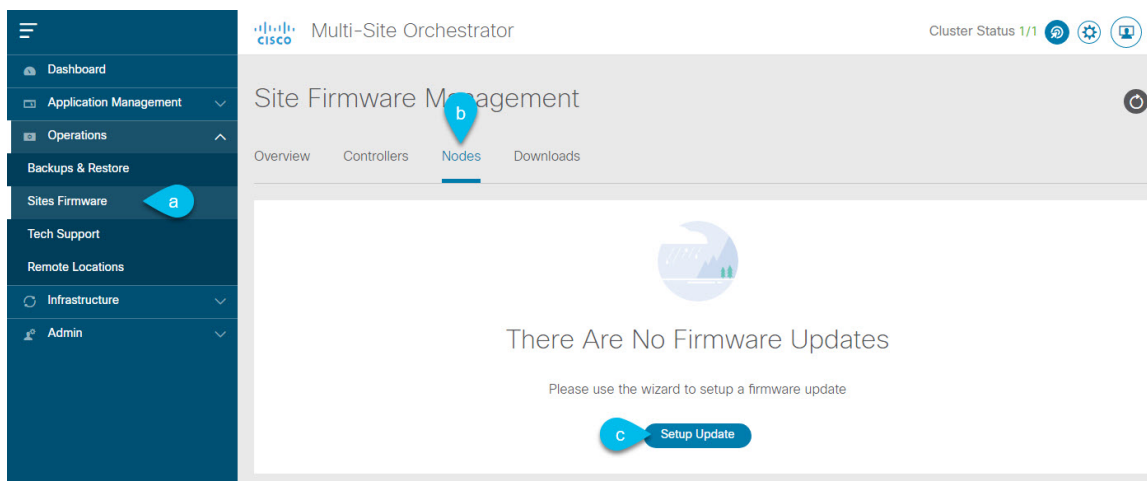
Upgrading Nodes

This section describes how to set up a software upgrade for your sites' switch nodes.

Procedure

Step 1 Log in to your Nexus Dashboard Orchestrator.

Step 2 Set up switch nodes upgrade.



- a) From the left navigation pane, select **Operations** > **Sites Firmware**.
- b) In the main window, select the **Nodes** tab.
- c) Click **Setup Update** tab.

If you have previously set up one or more updates, click the **Setup Update** button in the top right of the main pane instead.

The **Setup Node Firmware Update** screen opens.

Step 3 Provide the upgrade details.

- a) Provide the **Name**.

This is the upgrade policy name you will be able to use to track the upgrade progress at any time.

- b) Click **Select Nodes**.

The **Select Nodes** window opens.

- c) Select a site, then select the switch nodes in that site and click **Add and Close**.

You can add switch nodes from a single site at a time. You will repeat this step if you want to add switches from other sites.

- d) Repeat the previous substep for nodes in other sites
- e) Click **Next** to proceed.

Step 4 In the **Version Selection** screen, select the firmware version and click **Next**.

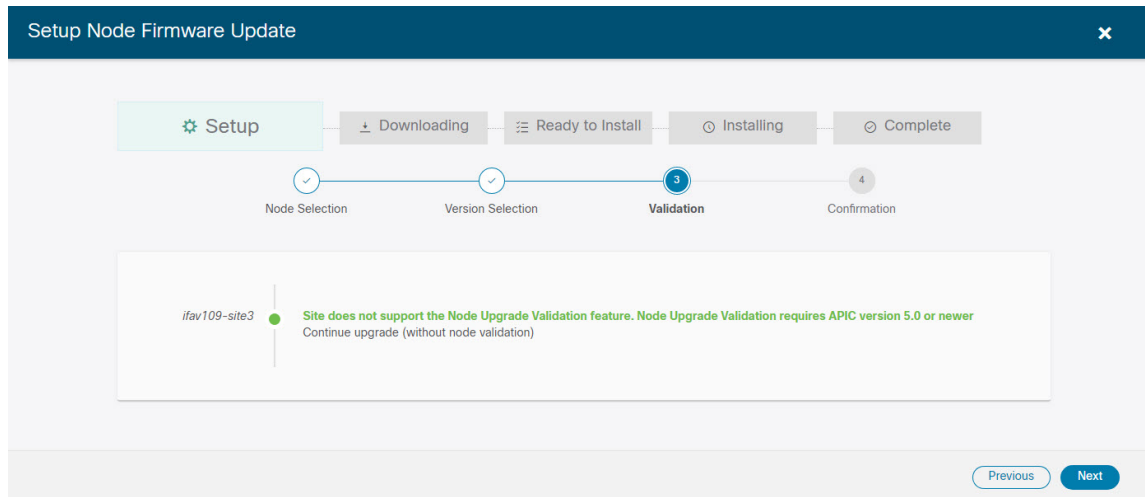
The firmware must be downloaded to the sites before it becomes available here. If the download you set up in previous section has completed successfully but the image is still not available here, close the **Setup Site Firmware Update** screen, navigate back to **Operations** > **Sites Firmware** > **Nodes** tab, and click the **Refresh** button to reload the latest information available for the sites; then restart the upgrade steps.

Step 5 In the **Validation** screen, ensure that there are no faults raised, then click **Next**.

Ensure that there are no faults and review any additional information that may affect your upgrade:

Note

Sites running releases prior to Release 5.0(1) do not support node validation, so we recommend checking for any switch node faults in the site's APIC prior to starting the upgrade from NDO.



Step 6 In the **Confirmation** screen, review the information and click **Submit**.

This will trigger image to be pre-downloaded to all the nodes you have selected. After the download completes, the screen will transition to **Ready to Install** and you can proceed to the next step.

Step 7 (Optional) Change **Advanced Settings**.

Note

Review the guidelines, recommendations, and limitations for the Cisco APIC upgrade process described in the [Upgrading and Downgrading the Cisco APIC and Switch Software](#) of the *Cisco APIC Installation, Upgrade, and Downgrade Guide* before making changes to the advanced options.

In the **Ready to Install** screen, you can open the **Advanced Settings** menu for additional options:

- **Ignore Compatibility Check**—by default, the option is set to **No** and compatibility check is enabled and verifies if an upgrade path from the currently-running version of the system to a specified newer version is supported.
If you choose to ignore the compatibility check feature, you run the risk of making an unsupported upgrade to your system, which could result in your system going to an unavailable state.
- **Graceful Check**—by default, the option is set to **No** and the upgrade process will not put any of the switches into Graceful Insertion and Removal (GIR) mode before performing the upgrade.
You can choose to enable this option to bring down the node gracefully (using GIR) while performing the upgrade so that the upgrade will have reduced traffic loss.
- **Run Mode**—by default, the option is set to **Continue on Failure** and if a node upgrade fails, the process proceeds to the next node. Alternatively, you can set this option to **Pause on Failure** to halt upgrade process if any one of the node upgrades fails.

Step 8 Remove any nodes marked as **Failed** from the upgrade.

The upgrade cannot proceed if the upgrade policy contains one or more nodes that failed to download the firmware. You can mouse over the `Failed` status for more information and reason for failure.

To remove the nodes from the upgrade, click **Edit Update Details** link in the **Ready to Install** screen.

Step 9 Click **Install** to start the upgrade.

If NDO to site connectivity is lost during the upgrade process, the GUI will display the last known status of the upgrade prior to loss of connectivity. Once connectivity is re-established, the upgrade status will be refreshed. You can perform a manual refresh after connectivity loss by clicking the **Refresh** button in the top right of the main pane.



CHAPTER 8

Tech Support

- [Tech Support and System Logs, on page 123](#)
- [Downloading System Logs, on page 124](#)
- [Streaming System Logs to External Analyzer, on page 124](#)

Tech Support and System Logs

Nexus Dashboard Orchestrator system logging is automatically enabled when you first deploy the Orchestrator cluster and captures the events and faults that occur in the environment.

You can choose to download the logs at any time or stream them to an external log analyzer, such as Splunk, if you want to use additional tools to quickly parse, view, and respond to important events without a delay.

Starting with Release 3.3(1), the tech support logs are split into two parts:

- Original database backup files containing the same information as in prior releases
- JSON-based database backup for ease of readability

Within each backup archive, you will find the following contents:

- `x.x.x.x`—one or more files in `x.x.x.x` format for container logs available at the time of the backup.
- `msc-backup-<date>_temp`—Original database backup containing the same information as previous releases.
- `msc-db-json-<date>_temp`—Backup contents in JSON format.

For example:

```
msc_anpEpgRels.json
msc_anpExtEpgRels.json
msc_asyncExecutionStatus.json
msc_audit.json
msc_backup-versions.json
msc_backupRecords.json
msc_ca-cert.json
msc_cloudSecStatus.json
msc_consistency.json
...
```

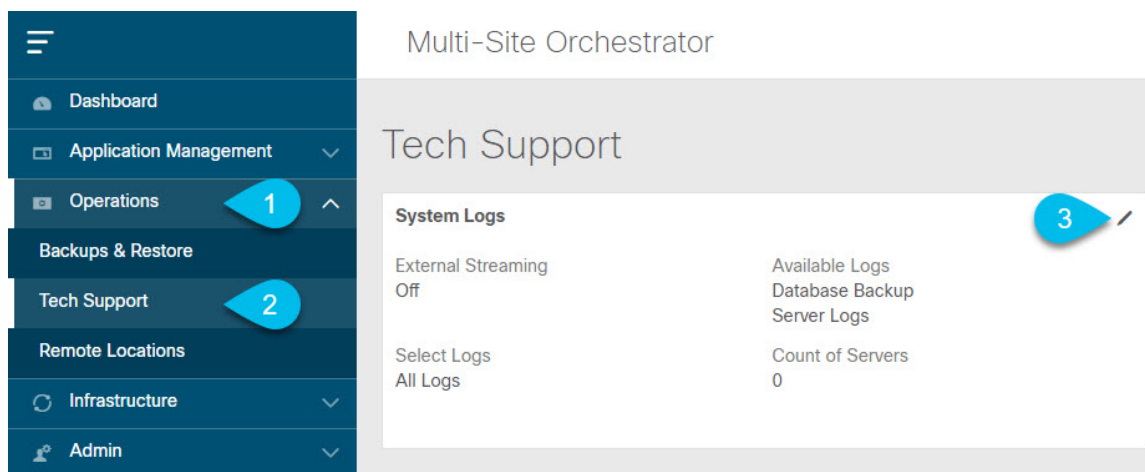
Downloading System Logs

This section describes how to generate a troubleshooting report and infrastructure logs file for all the schemas, sites, tenants, and users that are managed by Nexus Dashboard Orchestrator.

Procedure

Step 1 Log in to your Nexus Dashboard Orchestrator GUI.

Step 2 Open the **System Logs** screen.



- a) In the main menu, select **Operations > Tech Support**.
- b) In the top right corner of the **System Logs** frame, click the edit button.

Step 3 Click **Download** to download the logs.

An archive will be downloaded to your system. Containing all the information as described in the first section of this chapter.

Streaming System Logs to External Analyzer

Nexus Dashboard Orchestrator allows you to send the Orchestrator logs to an external log analyzer tool in real time. By streaming any events as they are generated, you can use the additional tools to quickly parse, view, and respond to important events without a delay.

This section describes how to enable Nexus Dashboard Orchestrator to stream its logs to an external analyzer tool, such as Splunk or syslog.

Before you begin

- This release supports only Splunk and `syslog` as external log analyzer.
- This release supports `syslog` only for Nexus Dashboard Orchestrator in Nexus Dashboard deployments.

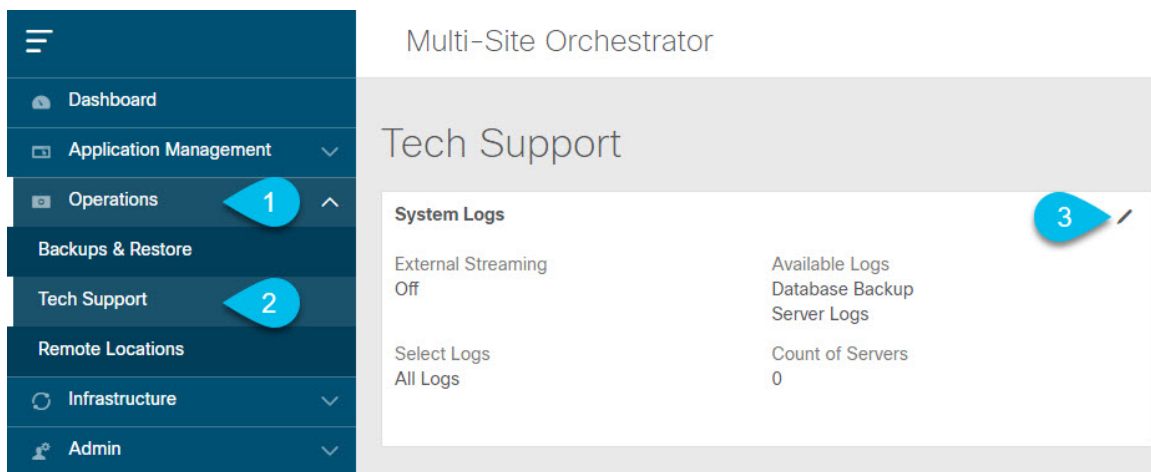
- This release supports up to 5 external servers.
- If using Splunk, set up and configure the log analyzer service provider.
For detailed instructions on how to configure an external log analyzer, consult its documentation.
- If using Splunk, obtain an authentication token for the service provider.

Obtaining an authentication token for Splunk service is detailed in the Splunk documentation, but in short, you can get the authentication token by logging into the Splunk server, selecting **Settings > Data Inputs > HTTP Event Collector**, and clicking **New Token**.

Procedure

Step 1 Log in to your Nexus Dashboard Orchestrator GUI.

Step 2 Open the **System Logs** screen.



- In the main menu, select **Operations > Tech Support**.
- In the top right corner of the **System Logs** frame, click the edit button.

Step 3 In the **System Logs** window, enable external streaming and add a server.

- Enable the **External Streaming** knob.
- Choose whether you want to stream **All Logs** or just the **Audit Logs**.
- Click **Add Server** to add an external log analyzer server.

Step 4

Add a Splunk server.

If you do not plan to use Splunk service, skip this step.

- Choose `splunk` for the server type.
- Choose the protocol.
- Provide the server name or IP address, port, and the authentication token you obtained from the Splunk service.

Obtaining an authentication token for Splunk service is detailed in the Splunk documentation, but in short, you can get the authentication token by logging into the Splunk server, selecting **Settings > Data Inputs > HTTP Event Collector**, and clicking **New Token**.

- Click the checkmark icon to finish adding the server.

Step 5 Add a `syslog` server.

If you do not plan to use `syslog`, skip this step.

The screenshot shows a configuration form for adding a syslog server. It includes the following fields and callouts:

- Select Service:** A dropdown menu with 'syslog' selected. A blue callout bubble with the number '1' points to the dropdown arrow.
- Protocol:** Two buttons, 'TCP' and 'UDP'. 'UDP' is selected. A blue callout bubble with the number '2' points to the 'UDP' button.
- * Host:** A text input field containing '10.195.223.220'. A vertical blue line is positioned to the right of this field.
- * Port:** A text input field containing '514'. A blue callout bubble with the number '3' points to the vertical blue line.
- Severity:** A dropdown menu with 'Warning' selected. A blue callout bubble with the number '4' points to the checkmark icon in the top right corner of the form.

- Choose `syslog` for the server type.
- Choose the protocol.
- Provide the server name or IP address, port number, and the severity level of the log messages to stream.
- Click the checkmark icon to finish adding the server.

Step 6 Repeat the steps if you want to add multiple servers.

This release supports up to 5 external servers.

Step 7 Click **Save** to save the changes.

System Logs

×

Download Logs

Download

External Streaming

☐

Select Logs

All Logs

Audit Logs

* Logging Servers ⓘ

Server Type	Protocol	Host	Port	
splunk	http	10.30.11.69	8088	✕
syslog	tcp	10.195.223.220	514	✕

+

 Add Server

SAVE



PART III

Infrastructure Management

- [System Configuration, on page 131](#)
- [Preparing Cisco APIC Sites, on page 133](#)
- [Adding and Deleting Sites, on page 141](#)
- [Configuring Infra General Settings, on page 147](#)
- [Configuring Infra for Cisco APIC Sites, on page 153](#)
- [Configuring Infra for Cisco Cloud Network Controller Sites, on page 161](#)
- [Deploying Infra Configuration for ACI Sites, on page 167](#)
- [CloudSec Encryption, on page 173](#)



CHAPTER 9

System Configuration

- [System Configuration Settings, on page 131](#)
- [System Alias and Banner, on page 131](#)

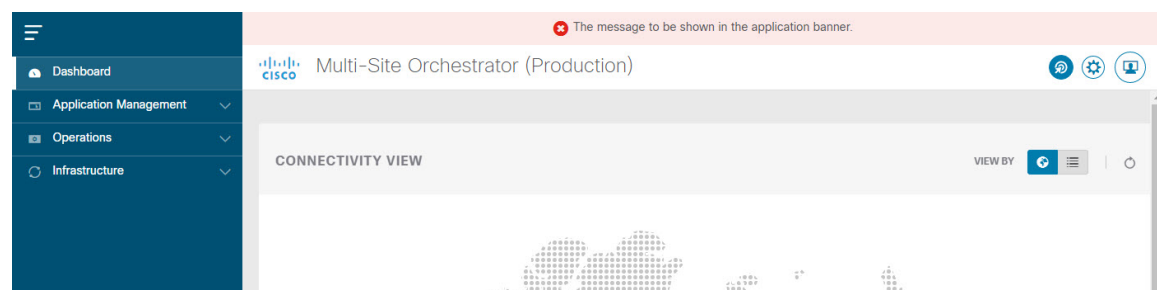
System Configuration Settings

There is a number of global system settings that are available under **Admin > System Configuration**, which you can configure for your Multi-Site Orchestrator as described in the following sections.

System Alias and Banner

This section describes how to configure an alias for your Nexus Dashboard Orchestrator as well as enable a custom GUI-wide banner to be displayed at the top of your screen, as shown in the following figure.

Figure 8: System Banner Display



Procedure

- Step 1** Log in to your Orchestrator.
- Step 2** From the left navigation pane, select **Admin > System Configuration**.
- Step 3** Click the **Edit** icon to the right of the **System Alias & Banners** area.
This opens the **System Alias & Banners** settings window.
- Step 4** In the **Alias** field, specify the system alias.

- Step 5** Choose whether you want to enable the GUI banner.
- Step 6** If you enable the banner, you must provide the message that will be displayed on it.
- Step 7** If you enable the banner, you must choose the severity, or color, for the banner.
- Step 8** Click **Save** to save the changes.
-



CHAPTER 10

Preparing Cisco APIC Sites

- [Pod Profile and Policy Group, on page 133](#)
- [Configuring Fabric Access Policies for All APIC Sites, on page 134](#)
- [Configuring Sites That Contain Remote Leaf Switches, on page 137](#)
- [Cisco Mini ACI Fabrics, on page 138](#)

Pod Profile and Policy Group

In each site's APIC, you must have one Pod profile with a Pod policy group. If your site does not have a Pod policy group you must create one. Typically, these settings will already exist as you will have configured them when you first deployed the fabric.

Procedure

- Step 1** Log in to the site's APIC GUI.
- Step 2** Check that the Pod profile contains a Pod policy group.
Navigate to **Fabric > Fabric Policies > Pods > Profiles > Pod Profile default**.
- Step 3** If necessary, create a Pod policy group.
a) Navigate to **Fabric > Fabric Policies > Pods > Policy Groups**.
b) Right-click **Policy Groups** and select **Create Pod Policy Group**.
c) Enter the appropriate information and click **Submit**.
- Step 4** Assign the new Pod policy group to the default Pod profile.
a) Navigate to **Fabric > Fabric Policies > Pods > Profiles > Pod Profile default**
b) Select the default profile.
c) Choose the new pod policy group and click **Update**.
-

Configuring Fabric Access Policies for All APIC Sites

Before your APIC fabrics can be added to and managed by the Nexus Dashboard Orchestrator, there is a number of fabric-specific access policies that you must configure on each site.

Configuring Fabric Access Global Policies

This section describes the global fabric access policy configurations that must be created for each APIC site before it can be added to and managed by the Nexus Dashboard Orchestrator.

Procedure

Step 1 Log in directly to the site's APIC GUI.

Step 2 From the main navigation menu, select **Fabric > Access Policies**.

You must configure a number of fabric policies before the site can be added to the Nexus Dashboard Orchestrator. From the APIC's perspective, this is something you do just like you would if you were connecting a bare-metal host, where you would configure domains, AEPs, policy groups, and interface selectors; you must configure the same options for connecting the spine switch interfaces to the inter-site network for all the sites that will be part of the same Multi-Site domain.

Step 3 Specify the VLAN pool.

The first thing you configure is the VLAN pool. We use Layer 3 sub-interfaces tagging traffic with VLAN-4 to connect the spine switches to the inter-site network.

- a) In the left navigation tree, browse to **Pools > VLAN**.
- b) Right-click the **VLAN** category and choose **Create VLAN Pool**.

In the **Create VLAN Pool** window, specify the following:

- For the **Name** field, specify the name for the VLAN pool, for example `msite`.
- For **Allocation Mode**, specify `Static Allocation`.
- And for the **Encap Blocks**, specify just the single VLAN 4. You can specify a single VLAN by entering the same number in both **Range** fields.

Step 4 Configure Attachable Access Entity Profiles (AEP).

- a) In the left navigation tree, browse to **Global Policies > Attachable Access Entity Profiles**.
- b) Right-click the **Attachable Access Entity Profiles** category and choose **Create Attachable Access Entity Profiles**.

In the **Create Attachable Access Entity Profiles** window, specify the name for the AEP, for example `msite-aep`.

- c) Click **Next** and **Submit**

No additional changes, such as interfaces, are required.

Step 5 Configure domain.

The domain you configure is what you will select from the Nexus Dashboard Orchestrator when adding this site.

- a) In the left navigation tree, browse to **Physical and External Domains > External Routed Domains**.
- b) Right-click the **External Routed Domains** category and choose **Create Layer 3 Domain**.

In the **Create Layer 3 Domain** window, specify the following:

- For the **Name** field, specify the name the domain, for example `msite-13`.
- For **Associated Attachable Entity Profile**, select the AEP you created in Step 4.
- For the **VLAN Pool**, select the VLAN pool you created in Step 3.

- c) Click **Submit**.

No additional changes, such as security domains, are required.

What to do next

After you have configured the global access policies, you must still add interfaces policies as described in [Configuring Fabric Access Interface Policies, on page 135](#).

Configuring Fabric Access Interface Policies

This section describes the fabric access interface configurations that must be done for the Nexus Dashboard Orchestrator on each APIC site.

Before you begin

You must have configured the global fabric access policies, such as VLAN Pool, AEP, and domain, in the site's APIC, as described in [Configuring Fabric Access Global Policies, on page 134](#).

Procedure

Step 1 Log in directly to the site's APIC GUI.

Step 2 From the main navigation menu, select **Fabric > Access Policies**.

In addition to the VLAN, AEP, and domain you have configured in previous section, you must also create the interface policies for the fabric's spine switch interfaces that connect to the Inter-Site Network (ISN).

Step 3 Configure a spine policy group.

- a) In the left navigation tree, browse to **Interface Policies > Policy Groups > Spine Policy Groups**.
This is similar to how you would add a bare-metal server, except instead of a Leaf Policy Group, you are creating a Spine Policy Group.
- b) Right-click the **Spine Policy Groups** category and choose **Create Spine Access Port Policy Group**.

In the **Create Spine Access Port Policy Group** window, specify the following:

- For the **Name** field, specify the name for the policy group, for example `Spine1-PolGrp`.
- For the **Link Level Policy** field, specify the link policy used between your spine switch and the ISN.

- For **CDP Policy**, choose whether you want to enable CDP.
- For the **Attached Entity Profile**, select the AEP you have configured in previous section, for example `msite-aep`.

c) Click **Submit**.

No additional changes, such as security domains, are required.

Step 4 Configure a spine profile.

- a) In the left navigation tree, browse to **Interface Policies > Profiles > Spine Profiles**.
- b) Right-click the **Spine Profiles** category and choose **Create Spine Interface Profile**.

In the **Create Spine Interface Profile** window, specify the following:

- For the **Name** field, specify the name for the profile, for example `Spine1-ISN`.
- For **Interface Selectors**, click the + sign to add the port on the spine switch that connects to the ISN. Then in the **Create Spine Access Port Selector** window, provide the following:
 - For the **Name** field, specify the name for the port selector, for example `Spine1-ISN`.
 - For the **Interface IDs**, specify the switch port that connects to the ISN, for example `5/32`.
 - For the **Interface Policy Group**, choose the policy group you created in the previous step, for example `Spine1-PolGrp`.

Then click **OK** to save the port selector.

c) Click **Submit** to save the spine interface profile.

Step 5 Configure a spine switch selector policy.

- a) In the left navigation tree, browse to **Switch Policies > Profiles > Spine Profiles**.
- b) Right-click the **Spine Profiles** category and choose **Create Spine Profile**.

In the **Create Spine Profile** window, specify the following:

- For the **Name** field, specify the name for the profile, for example `Spine1`.
- For **Spine Selectors**, click the + to add the spine and provide the following:
 - For the **Name** field, specify the name for the selector, for example `Spine1`.
 - For the **Blocks** field, specify the spine node, for example `201`.

- c) Click **Update** to save the selector.
- d) Click **Next** to proceed to the next screen.
- e) Select the interface profile you have created in the previous step

For example `Spine1-ISN`.

f) Click **Finish** to save the spine profile.

Configuring Sites That Contain Remote Leaf Switches

Multi-Site architecture supports APIC sites with Remote Leaf switches. The following sections describe guidelines, limitations, and configuration steps required to allow Nexus Dashboard Orchestrator to manage these sites.

Remote Leaf Guidelines and Limitations

If you want to add an APIC site with a Remote Leaf to be managed by the Nexus Dashboard Orchestrator, the following restrictions apply:

- You must upgrade your Cisco APIC to Release 4.2(4) or later.
- Only physical Remote Leaf switches are supported in this release
- Only -EX and -FX or later switches are supported as Remote Leaf switches for use with Multi-Site
- Remote Leaf is not supported with back-to-back connected sites without IPN switches
- Remote Leaf switches in one site cannot use another site's L3Out
- Stretching a bridge domain between one site and a Remote Leaf in another site is not supported

You must also perform the following tasks before the site can be added to and managed by the Nexus Dashboard Orchestrator:

- You must enable Remote Leaf direct communication and configure routable subnets directly in the site's APIC, as described in the following sections.
- You must add the routable IP addresses of Cisco APIC nodes in the DHCP-Relay configuration applied on the interfaces of the Layer 3 routers connecting to the Remote Leaf switches.

The routable IP address of each APIC node is listed in the **Routable IP** field of the **System > Controllers > <controller-name>** screen of the APIC GUI.

Configuring Routable Subnets for Remote Leaf Switches

Before you can add a site that contains one or more Remote Leaf switches to the Nexus Dashboard Orchestrator, you must configure routable subnets for the pod with which the Remote Leaf nodes are associated.

Procedure

-
- | | |
|---------------|---|
| Step 1 | Log in directly to the site's APIC GUI. |
| Step 2 | From the menu bar, select Fabric > Inventory . |
| Step 3 | In the Navigation pane, click Pod Fabric Setup Policy . |
| Step 4 | In the main pane, double-click the pod where you want to configure the subnets. |
| Step 5 | In the Routable Subnets area, click the + sign to add a subnet. |
| Step 6 | Enter the IP and Reserve Address Count , set the state to Active or Inactive , then click Update to save the subnet. |

When configuring routable subnets, you must provide a netmask between /22 and /29.

Step 7 Click **Submit** to save the configuration.

Enabling Direct Communication for Remote Leaf Switches

Before you can add a site that contains one or more Remote Leaf switches to the Nexus Dashboard Orchestrator, you must configure direct remote leaf communication for that site. Additional information about remote leaf direct communication feature is available in the *Cisco APIC Layer 3 Networking Configuration Guide*. This section outlines the steps and guidelines specific to the integration with Multi-Site.



Note Once you enable Remote Leaf switch direct communication, the switches will function in the new mode only

Procedure

Step 1 Log in directly to the site's APIC.

Step 2 Enable direct traffic forwarding for Remote Leaf switches.

- a) From the menu bar, navigate to **System > System Settings**.
- b) From the left side bar, select **Fabric Wide Setting**.
- c) Check the **Enable Remote Leaf Direct Traffic Forwarding** checkbox.

Note

You cannot disable this option after you enable it.

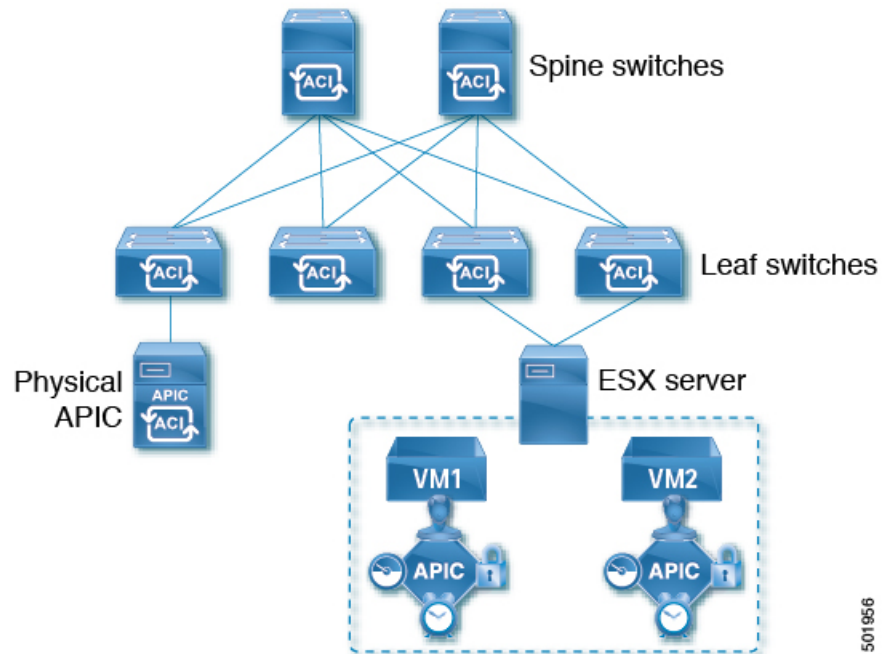
- d) Click **Submit** to save the changes.

Cisco Mini ACI Fabrics

Cisco Multi-Site supports Cisco Mini ACI fabrics as typical on-premises sites without requiring any additional configuration. This section provides a brief overview of Mini ACI fabrics, detailed info on deploying and configuring this type of fabrics is available in [Cisco Mini ACI Fabric and Virtual APICs](#).

Cisco ACI, Release 4.0(1) introduced Mini ACI Fabric for small scale deployment. Mini ACI fabric works with Cisco APIC cluster consisting of one physical APIC and two virtual APICs (vAPIC) running in virtual machines. This reduces the physical footprint and cost of the APIC cluster, allowing ACI fabric to be deployed in scenarios with limited rack space or initial budget, such as a colocation facility or a single-room data center, where a full-scale ACI installations may not be practical due to physical footprint or initial cost.

The following diagram shows an example of a mini Cisco ACI fabric with a physical APIC and two virtual APICs (vAPICs):

Figure 9: Cisco Mini ACI Fabric

501956



CHAPTER 11

Adding and Deleting Sites

- [Cisco NDO and APIC Interoperability Support, on page 141](#)
- [Adding Cisco ACI Sites, on page 143](#)
- [Removing Sites, on page 144](#)
- [Cross Launch to Fabric Controllers, on page 145](#)

Cisco NDO and APIC Interoperability Support

Cisco Nexus Dashboard Orchestrator (NDO) does not require a specific version of APIC to be running in all sites. The APIC clusters in each site as well as the NDO itself can be upgraded independently of each other and run in mixed operation mode as long as the fabric can be on-boarded to the Nexus Dashboard where the Nexus Dashboard Orchestrator service is installed. As such, we recommend that you always upgrade to the latest release of the Nexus Dashboard Orchestrator.

However, keep in mind that if you upgrade the NDO before upgrading the APIC clusters in one or more sites, some of the new NDO features may not yet be supported by an earlier APIC release. In that case a check is performed on each template to ensure that every configured option is supported by the target sites.

The check is performed when you save a template or deploy a template. If the template is already assigned to a site, any unsupported configuration options will not be saved; if the template is not yet assigned, you will be able to assign it to a site, but not be able to save or deploy the schema if it contains configuration unsupported by that site.

In case an unsupported configuration is detected, an error message will show, for example: This APIC site version `<site-version>` is not supported by NDO. The minimum version required for this `<feature>` is `<required-version>` or above.

The following table lists the features and the minimum required APIC release for each one:



Note While some of the following features are supported on earlier Cisco APIC releases, Release 4.2(4) is the earliest release that can be on-boarded to the Nexus Dashboard and managed by this release of Nexus Dashboard Orchestrator.

Feature	Minimum APIC Version
ACI Multi-Pod Support	Release 4.2(4)

Feature	Minimum APIC Version
Service Graphs (L4-L7 Services)	Release 4.2(4)
External EPGs	Release 4.2(4)
ACI Virtual Edge VMM Support	Release 4.2(4)
DHCP Support	Release 4.2(4)
Consistency Checker	Release 4.2(4)
vzAny	Release 4.2(4)
Host Based Routing	Release 4.2(4)
CloudSec Encryption	Release 4.2(4)
Layer 3 Multicast	Release 4.2(4)
MD5 Authentication for OSPF	Release 4.2(4)
EPG Preferred Group	Release 4.2(4)
Intersite L3Out	Release 4.2(4)
EPG QoS Priority	Release 4.2(4)
Contract QoS Priority	Release 4.2(4)
Single Sign-On (SSO)	Release 5.0(1)
Multicast Rendezvous Point (RP) Support	Release 5.0(1)
Transit Gateway (TGW) support for AWS and Azure Sites	Release 5.0(1)
SR-MPLS Support	Release 5.0(1)
Cloud LoadBalancer High Availability Port	Release 5.0(1)
Service Graphs (L4-L7 Services) with UDR	Release 5.0(2)
3rd Party Device Support in Cloud	Release 5.0(2)
Cloud Loadbalancer Target Attach Mode Feature	Release 5.1(1)
Support security and service insertion in Azure for non-ACI networks reachable through Express Route	Release 5.1(1)
CSR Private IP Support	Release 5.1(1)
Extend ACI policy model and automation for Cloud native services in Azure	Release 5.1(1)

Feature	Minimum APIC Version
Flexible segmentation through multiple VRF support within a single VNET for Azure	Release 5.1(1)
Private Link automation for Azure PaaS and third-party services	Release 5.1(1)
Openshift 4.3 IPI on Azure with ACI-CNI	Release 5.1(1)
Cloud Site Underlay Configuration	Release 5.2(1)

Adding Cisco ACI Sites

This section describes how to add a Cisco APIC or Cloud Network Controller site using the Nexus Dashboard GUI and then enable that site to be managed by Nexus Dashboard Orchestrator.

Before you begin

- If you are adding on-premises ACI site, you must have completed the site-specific configurations in each site's APIC, as described in previous sections in this chapter.
- You must ensure that the site(s) you are adding are running Release 4.2(4) or later.

Procedure

Step 1 Log in to your Nexus Dashboard and open the **Admin Console**.

Step 2 From the left navigation menu, choose **Sites** and click **Add Site..**

Step 3 Provide site information.

- a) For **Site Type**, select **ACI** or **Cloud Network Controller** depending on the type of ACI fabric you are adding.
- b) Provide the controller information.

- You need to provide the **Host Name/IP Address**, **User Name**, and **Password**. for the APIC controller currently managing your ACI fabrics.

Note

For APIC fabrics, if you will use the site with Nexus Dashboard Orchestrator service only, you can provide either the in-band or out-of-band IP address of the APIC. If you will use the site with Nexus Dashboard Insights as well, you must provide the in-band IP address.

- For on-premises ACI sites managed by Cisco APIC, if you plan to use this site with Day-2 Operations applications such as Nexus Insights, you must also provide the **In-Band EPG** name used to connect the Nexus Dashboard to the fabric you are adding. Otherwise, if you will use this site with Nexus Dashboard Orchestrator only, you can leave this field blank.
- For Cloud Network Controller sites, **Enable Proxy** if your cloud site is reachable via a proxy.

Proxy must be already configured in your Nexus Dashboard's cluster settings. If the proxy is reachable via management network, a static management network route must also be added for the proxy IP address. For more information about proxy and route configuration, see [Nexus Dashboard User Guide](#) for your release.

- c) Click **Save** to finish adding the site.

At this time, the sites will be available in the Nexus Dashboard, but you still need to enable them for Nexus Dashboard Orchestrator management as described in the following steps.

Step 4 Repeat the previous steps for any additional ACI or Cloud Network Controller sites.

Step 5 From the Nexus Dashboard's **Services** page, open the Nexus Dashboard Orchestrator service.
You will be automatically logged in using the Nexus Dashboard user's credentials.

Step 6 In the Nexus Dashboard Orchestrator GUI, manage the sites.

- a) From the left navigation menu, select **Sites**.
- b) In the main pane, change the **State** from *Unmanaged* to *Managed* for each fabric that you want the NDO to manage.

When managing the sites, you must provide a unique site ID for each site.

Note

Ensure that ACI site names are limited to 125 characters or less to avoid any issues when enabling orchestration.

Removing Sites

This section describes how to disable site management for one or more sites using the Nexus Dashboard Orchestrator GUI. The sites will remain present in the Nexus Dashboard.

Before you begin

You must ensure that all templates associated with the site you want to remove are not deployed.

Procedure

Step 1 Open the Nexus Dashboard Orchestrator GUI.

You can open the NDO service from the Nexus Dashboard's **Service Catalog**. You will be automatically logged in using the Nexus Dashboard user's credentials.

Step 2 Remove the site from all templates.

You must remove the site from all templates with which it is associated before you can unmanage the site and remove it from your Nexus Dashboard.

- a) Navigate to **Application Management > Schemas**.
- b) Click a schema that contains one or more templates associated with the site.
- c) In the left sidebar's **Sites** area, select a template associated with the site, click the options menu (...) next to the template, and choose **Undeploy Template**.

This will remove configurations that were deployed using this template to this site.

Note

For non-stretched templates, you can choose to preserve the configuration by selecting **Dissociate Template** instead of **Undeploy Template**, but you must undeploy any stretched templates.

- d) Repeat this step for all templates associated with the site that you want to unmanage in this and all other schemas.

Step 3

Remove the site's underlay configuration.

- a) From the left navigation menu, select **Infrastructure > Site Connectivity**.
- b) In the main pane, click **Configure**.
- c) In the left sidebar, select the site you want to unmanage.
- d) In right sidebar's **Inter-Site Connectivity** tab, disable the **Multi-Site** checkbox.
- e) Click **Deploy** to deploy the changes to the site.

Step 4

In the Nexus Dashboard Orchestrator GUI, disable the sites.

- a) From the left navigation menu, select **Infrastructure > Sites**.
- b) In the main pane, change the **State** from *Managed* to *Unmanaged* for the site that you want to unmanage.

Note

If the site is associated with one or more deployed templates, you will not be able to change its state to *Unmanaged* until you undeploy those templates, as described in the previous step.

Step 5

Delete the site from Nexus Dashboard.

If you no longer want to manage this site or use it with any other applications, you can delete the site from the Nexus Dashboard as well.

Note

Note that the site must not be currently in use by any of the services installed in your Nexus Dashboard cluster.

- a) In the top navigation bar, click the **Home** icon to return to the Nexus Dashboard GUI.
- b) From the left navigation menu of the Nexus Dashboard GUI, select **Sites**.
- c) Select one or more sites you want to delete.
- d) In the top right of the main pane, select **Actions > Delete Site**.
- e) Provide the site's login information and click **OK**.

The site will be removed from the Nexus Dashboard.

Cross Launch to Fabric Controllers

Nexus Dashboard Orchestrator currently supports a number of configuration options for each type of fabrics. For many additional configuration options, you may need to log in directly into the fabric's controller.

You can cross launch into the specific site controller's GUI from the NDO's **Infrastructure > Sites** screen by selecting the actions (. . .) menu next to the site and clicking **Open in user interface**. Note that cross-launch works with out-of-band (OOB) management IP of the fabric.

If the same user is configured in Nexus Dashboard and the fabric, you will be logged in automatically into the fabric's controller using the same log in information as the Nexus Dashboard user. For consistency, we recommend configuring remote authentication with common users across Nexus Dashboard and the fabrics.

CHAPTER 12

Configuring Infra General Settings

- [Infra Configuration Dashboard, on page 147](#)
- [Partial Mesh Intersite Connectivity, on page 148](#)
- [Configuring Infra: General Settings, on page 149](#)

Infra Configuration Dashboard

The **Infra Configuration** page displays an overview of all sites and inter-site connectivity in your Nexus Dashboard Orchestrator deployment and contains the following information:

Figure 10: Infra Configuration Overview

The screenshot displays the Cisco Nexus Dashboard Orchestrator interface. The left sidebar contains navigation links: Dashboard, Sites, Application Management, Fabric Management, Operations, Infrastructure, System Configuration, Site Connectivity, and Integration. The main content area is titled 'Site Connectivity' and includes a 'Configure' button. It displays 'General Settings' for BGP peering type (full-mesh) and configuration (Keep Alive Interval: 60s, Hold Interval: 180s, BGP TTL: 16, State Interval: 300s, Graceful Restart: On, Maximum AS Limit: 0). Below this, it shows details for two sites: 'scale-ms11' (1 Pod, 1 Spine, ACI Multi-Site On, Cloudsec Encryption On, APIC Site ID 254, BGP Autonomous Sys Number 511, OSPF Area ID 0, OSPF Area Type regular, Overlay Multicast TEP 11.11.11.10, External Routed Domain uni/l3dom-L3dom) and 'Azsite1' (4 Regions, ACI Multi-Site On, APIC Site ID 21, BGP Autonomous Sys Number 65145). At the bottom, there is a table for 'Inter-Site Connections' with columns for Site Name, Deployment Status, Operational Status, Overlay Routing Status, and Tunnel Status. The table shows one connection 'onPrem2' with Deployment Status OK, Operational Status Fail, Overlay Routing Status 8 up 0 down 8, and Tunnel Status 4 up 0 down 4.

1. The **General Settings** tile displays information about BGP peering type and its configuration. This is described in detail in the next section.

2. The **On-Premises** tiles display information about every on-premises site that is part of your Multi-Site domain along with their number of Pods and spine switches, OSPF settings, and overlay IPs.

You can click on the **Pods** tile that displays the number of Pods in the site to show information about the Overlay Unicast TEP addresses of each Pod.

This is described in detail in [Configuring Infra for Cisco APIC Sites, on page 153](#).

3. The **Cloud** tiles display information about every cloud site that is part of your Multi-Site domain along with their number of regions and basic site information.

This is described in detail in [Configuring Infra for Cisco Cloud Network Controller Sites, on page 161](#).

4. You can click **Show Connectivity Status** to display intersite connectivity details for a specific site.
5. You can use the **Configure** button to navigate to the intersite connectivity configuration, which is described in detail in the following sections.

The following sections describe the steps necessary to configure the general fabric Infra settings. Fabric-specific requirements and procedures are described in the following chapters based on the specific type of fabric you are managing.

Before you proceed with Infra configuration, you must have configured and added the sites as described in previous sections.

In addition, any infrastructure changes such as adding and removing spine switches or spine node ID changes require a Nexus Dashboard Orchestrator fabric connectivity information refresh described in the [Refreshing Site Connectivity Information, on page 153](#) as part of the general Infra configuration procedures.

Partial Mesh Intersite Connectivity

In addition to full mesh connectivity where you configure intersite connectivity from every site managed by your Nexus Dashboard Orchestrator to every other site, this release also supports partial mesh configuration. In partial mesh configuration, you can manage sites in standalone mode with no intersite connectivity to any other site or limit the intersite configuration to only a subset of other sites in your Multi-Site domain.

Prior to Nexus Dashboard Orchestrator, Release 3.6(1), you could stretch templates between sites and refer to policies from other templates, which were deployed to other sites, even if the intersite connectivity between those sites was not configured, resulting in intended traffic flow between the sites to not work.

Beginning with release 3.6(1), the Orchestrator will allow you to stretch template and remote reference policies from other templates (deployed on other sites) between two or more sites only if the intersite connectivity between those sites is properly configured and deployed.

When configuring site infra for Cisco APIC and Cisco Cloud Network Controller sites as described in the following sections, for each site you can explicitly choose to which other sites infra connectivity will be established and provide that configuration information only.

Partial Mesh Connectivity Guidelines

When configuring partial mesh connectivity, consider the following guidelines:

- Partial mesh connectivity is supported between two cloud sites or a cloud and on-premises site.
Full mesh connectivity is automatically established between all on-premises sites.
- Partial mesh connectivity is supported using BGP-EVPN or BGP-IPv4 protocols.

Note however that stretching a template is allowed only for sites that are connected using BGP-EVPN protocol. If you are using BGP-IPv4 to connect two or more sites, any template assigned to any of those sites can be deployed to one site only.

Configuring Infra: General Settings

This section describes how to configure general Infra settings for all the sites.



Note Some of the following settings apply to all sites, while others are required for specific type of sites (for example, Cloud Network Controller sites). Ensure that you complete all the required configurations in infra general settings before proceeding to the site-local settings specific to each site.

Procedure

- Step 1** Log in to the Cisco Nexus Dashboard Orchestrator GUI.
- Step 2** In the left navigation menu, select **Infrastructure > Site Connectivity**.
- Step 3** In the main pane, click **Configure**.
- Step 4** In the left sidebar, select **General Settings**.
- Step 5** Provide **Control Plane Configuration**.
 - a) Select the **Control Plane Configuration** tab.
 - b) Choose **BGP Peering Type**.
 - **full-mesh**—All border gateway switches in each site will establish peer connectivity with remote sites' border gateway switches.

In **full-mesh** configuration, Nexus Dashboard Orchestrator uses the spine switches for ACI managed fabrics and border gateways for NDFC managed fabrics.
 - **route-reflector**—The route-reflector option allows you to specify one or more control-plane nodes to which each site establishes MP-BGP EVPN sessions. The use of route-reflector nodes avoids creating MP-BGP EVPN full mesh adjacencies between all the sites managed by NDO.

For ACI fabrics, the **route-reflector** option is effective only for fabrics that are part of the same BGP ASN.
 - c) In the **Keepalive Interval (Seconds)** field, enter the keep alive interval seconds.

We recommend keeping the default value.
 - d) In the **Hold Interval (Seconds)** field, enter the hold interval seconds.

We recommend keeping the default value.
 - e) In the **Stale Interval (Seconds)** field, enter stale interval seconds.

We recommend keeping the default value.
 - f) Choose whether you want to turn on the **Graceful Helper** option.
 - g) Provide the **Maximum AS Limit**.

We recommend keeping the default value.

- h) Provide the **BGP TTL Between Peers**.

We recommend keeping the default value.

- i) Provide the **OSPF Area ID**.

If you do not have any Cloud Network Controller sites, this field will not be present in the UI.

This is OSPF area ID used by cloud sites for on-premises IPN peering.

- j) (Optional) Enable **IANA Assigned Port** for CloudSec encryption.

By default, CloudSec uses a proprietary UDP port. This option allows you to configure CloudSec to use the official IANA-reserved port 8017 for CloudSec encryption between sites.

Note

The IANA-reserved port is supported for Cisco APIC sites running release 5.2(4) or later.

To change this setting, CloudSec must be disabled on all sites. If you want to enable IANA reserved port, but already have CloudSec encryption enabled for one or more of your sites, disable CloudSec for all sites, enable **IANA Reserve UDP Port** option, then re-enable CloudSec for the required sites.

For detailed information and steps for configuring CloudSec, see the "CloudSec Encryption" chapter of the [Nexus Dashboard Orchestrator Configuration Guide for ACI Fabrics](#).

Step 6 Provide the **IPN Devices** information.

If you do not plan to configure inter-site connectivity between on-premises and cloud sites, you can skip this step.

When you configure inter-site underlay connectivity between on-premises and cloud sites as described in later sections, you will need to select an on-premises IPN device which will establish connectivity to the cloud CSRs. These IPN devices must first be defined here before they are available in the on-premises site configuration screen, which is described in more detail in [Configuring Infra: On-Premises Site Settings, on page 154](#).

- a) Select the **On Premises IPSec Devices** tab.
- b) Click **+Add On-Premises IPSec Device**.
- c) Choose whether the device is **Unmanaged** or **Managed** and provide the device information.

This defines whether or not the device is directly managed by NDFC:

- For **Unmanaged** IPN devices, simply provide the **Name** and the **IP Address** of the device.

The IP address you provide will be used as the tunnel peer address from the cloud CSRs, not the IPN device's management IP address.

- For **Managed** IPN devices, choose the NDFC **Site** that contains the device and then the **Device** from that site.

Then choose the **Interface** on the device that is facing the Internet and provide the **Next Hop** IP address, which is the IP address of the gateway that is connecting to the Internet.

- d) Click the check mark icon to save the device information.
- e) Repeat this step for any additional IPN devices you want to add.

Step 7 Provide the **External Devices** information.

If you do not have any Cloud Network Controller sites, this tab will not be present in the UI.

If you do not have any Cloud Network Controller sites in your Multi-Site domain or you do not plan to configure connectivity between cloud sites and branch routers or other external devices, you can skip this step.

The following steps describe how to provide information about any branch routers or external devices to which you want to configure connectivity from your cloud sites.

- a) Select the **External Devices** tab.

This tab will only be available if you have at least one cloud site in your Multi-Site domain.

- b) Click **Add External Device**.

The **Add External Device** dialogue will open.

- c) Provide the **Name**, **IP Address**, and **BGP Autonomous System Number** for the device.

The IP address you provide will be used as the tunnel peer address from the Cloud Network Controller's CSRs, not the device's management IP address. The connectivity will be established over public Internet using IPSec.

- d) Click the check mark icon to save the device information.
- e) Repeat this step for any additional IPN devices you want to add.

After you have added all the external devices, ensure to complete the next step to provide the IPSec tunnel subnet pools from which the internal IP addresses will be allocated for these tunnels.

Step 8 Provide the **IPSec Tunnel Subnet Pools** information.

If you do not have any Cloud Network Controller sites, this tab will not be present in the UI.

There are two types of subnet pools that you can provide here:

- **External Subnet Pool**—used for connectivity between cloud site CSRs and other sites (cloud or on-premises).

These are large global subnet pools that are managed by Nexus Dashboard Orchestrator. The Orchestrator creates smaller subnets from these pools and allocates them to sites to be used for inter-site IPsec tunnels and external connectivity IPsec tunnels.

You must provide at least one external subnet pool if you want to enable external connectivity from one or more of your cloud sites.

- **Site-Specific Subnet Pool**—used for connectivity between cloud site CSRs and external devices.

These subnets can be defined when the external connectivity IPsec tunnels must be in a specific range. For example, where a specific subnet is already being used to allocate IP addresses to the external router and you want to continue using those subnets for IPsec tunnels for NDO and cloud sites. These subnets are not managed by the Orchestrator and each subnet is assigned to a site in its entirety to be used locally for external connectivity IPsec tunnels.

If you do not provide any named subnet pools but still configure connectivity between cloud site's CSRs and external devices, the external subnet pool will be used for IP allocation.

Note

The minimum mask length for both subnet pools is /24.

To add one or more **External Subnet Pools**:

- a) Select the **IPSec Tunnel Subnet Pools** tab.
- b) In the **External Subnet Pool** area, click **+Add IP Address** to add one or more external subnet pools.

This subnet will be used to address the IPsec tunnel interfaces and loopbacks of the Cloud Routers used for on-premises connectivity, which you previously configured in the Cloud Network Controller for inter-site connectivity in earlier Nexus Dashboard Orchestrator releases.

The subnets must not overlap with other on-premises TEP pools, should not begin with 0.x.x.x or 0.0.x.x, and should have a network mask between /16 and /24, for example 30.29.0.0/16.

- c) Click the check mark icon to save the subnet information.
- d) Repeat these substeps for any additional subnet pools you want to add.

To add one or more **Site-Specific Subnet Pools**:

- a) Select the **IPSec Tunnel Subnet Pools** tab.
- b) In the **Site-Specific Subnet Pools** area, click **+Add IP Address** to add one or more external subnet pools.

The **Add Named Subnet Pool** dialogue will open.

- c) Provide the subnet **Name**.

You will be able to use the subnet pool's name to choose the pool from which to allocate the IP addresses later on.

- d) Click **+Add IP Address** to add one or more subnet pools.

The subnets must have a network mask between /16 and /24 and not begin with 0.x.x.x or 0.0.x.x, for example 30.29.0.0/16.

- e) Click the check mark icon to save the subnet information.

Repeat the steps if you want to add multiple subnets to the same named subnet pool.

- f) Click **Save** to save the named subnet pool.
- g) Repeat these substeps for any additional named subnet pools you want to add.

What to do next

After you have configured general infra settings, you must still provide additional information for site-specific configurations based on the type of sites (ACI, Cloud Network Controller, or NDFC) you are managing. Follow the instructions described in the following sections to provide site-specific infra configurations.



CHAPTER 13

Configuring Infra for Cisco APIC Sites

- [Refreshing Site Connectivity Information, on page 153](#)
- [Configuring Infra: On-Premises Site Settings, on page 154](#)
- [Configuring Infra: Pod Settings, on page 156](#)
- [Configuring Infra: Spine Switches, on page 157](#)

Refreshing Site Connectivity Information

Any infrastructure changes, such as adding and removing spines or changing spine node IDs, require a Multi-Site fabric connectivity site refresh. This section describes how to pull up-to-date connectivity information directly from each site's APIC.

Procedure

-
- | | |
|---------------|---|
| Step 1 | Log in to the Cisco Nexus Dashboard Orchestrator GUI. |
| Step 2 | In the left navigation menu, select Infrastructure > Site Connectivity . |
| Step 3 | In the top right of the main pane, click Configure . |
| Step 4 | In the left pane, under Sites , select a specific site. |
| Step 5 | In the main window, click the Refresh button to pull fabric information from the APIC. |
| Step 6 | (Optional) For on-premises sites, in the Confirmation dialog, check the box if you want to remove configuration for decommissioned spine switch nodes.

If you choose to enable this checkbox, all configuration info for any currently decommissioned spine switches will be removed from the database. |
| Step 7 | Finally, click Yes to confirm and load the connectivity information.

This will discover any new or removed spines and all site-related fabric connectivity will be re-imported from the APIC. |
-

Configuring Infra: On-Premises Site Settings

This section describes how to configure site-specific Infra settings for on-premises sites.

Procedure

-
- Step 1** Log in to the Cisco Nexus Dashboard Orchestrator GUI.
- Step 2** In the left navigation menu, select **Infrastructure > Site Connectivity**.
- Step 3** In the top right of the main pane, click **Configure**.
- Step 4** In the left pane, under **Sites**, select a specific on-premises site.
- Step 5** Provide the **Inter-Site Connectivity** information.
- In the right **<Site> Settings** pane, enable the **Multi-Site** knob.
This defines whether the overlay connectivity is established between this site and other sites.
 - (Optional) Enable the **CloudSec Encryption** knob encryption for the site.
CloudSec Encryption provides inter-site traffic encryption. The "Infrastructure Management" chapter in the [Cisco Multi-Site Configuration Guide](#) covers this feature in detail.
 - Specify the **Overlay Multicast TEP**.
This address is used for the inter-site L2 BUM and L3 multicast traffic. This IP address is deployed on all spine switches that are part of the same fabric, regardless of whether it is a single pod or Multi-Pod fabric.
This address should not be taken from the address space of the original fabric's `Infra` TEP pool or from the `0.x.x.x` range.
 - Specify the **BGP Autonomous System Number**.
 - (Optional) Specify the **BGP Password**.
 - Provide the **OSPF Area ID**.
The following settings are required if you are using OSPF protocol for underlay connectivity between the site and the IPN. If you plan to use BGP instead, you can skip this step. BGP underlay configuration is done at the port level, as described in [Configuring Infra: Spine Switches, on page 157](#).
 - Select the **OSPF Area Type** from the dropdown menu.
The following settings are required if you are using OSPF protocol for underlay connectivity between the site and the IPN. If you plan to use BGP instead, you can skip this step. BGP underlay configuration is done at the port level, as described in [Configuring Infra: Spine Switches, on page 157](#).
The OSPF area type can be one of the following:
 - `nssa`
 - `regular`
 - Configure OSPF policies for the site.
The following settings are required if you are using OSPF protocol for underlay connectivity between the site and the IPN. If you plan to use BGP instead, you can skip this step. BGP underlay configuration is done at the port level, as described in [Configuring Infra: Spine Switches, on page 157](#).

You can either click an existing policy (for example, `mso-ospf-policy-default`) to modify it or click **+Add Policy** to add a new OSPF policy. Then in the **Add/Update Policy** window, specify the following:

- In the **Policy Name** field, enter the policy name.
 - In the **Network Type** field, choose either `broadcast`, `point-to-point`, or `unspecified`.
The default is `broadcast`.
 - In the **Priority** field, enter the priority number.
The default is 1.
 - In the **Cost of Interface** field, enter the cost of interface.
The default is 0.
 - From the **Interface Controls** dropdown menu, choose one of the following:
 - **advertise-subnet**
 - **bfd**
 - **mtu-ignore**
 - **passive-participation**
 - In the **Hello Interval (Seconds)** field, enter the hello interval in seconds.
The default is 10.
 - In the **Dead Interval (Seconds)** field, enter the dead interval in seconds.
The default is 40.
 - In the **Retransmit Interval (Seconds)** field, enter the retransmit interval in seconds.
The default is 5.
 - In the **Transmit Delay (Seconds)** field, enter the transmit delay in seconds.
The default is 1.
- i) (Optional) From the **External Routed Domain** dropdown, select the domain you want to use.
Choose an external router domain that you have created in the Cisco APIC GUI. For more information, see the *Cisco APIC Layer 3 Networking Configuration Guide* specific to your APIC release.
- j) (Optional) Enable **SDA Connectivity** for the site.
If the site is connected to an SDA network, enable the **SDA Connectivity** knob and provide the **External Routed Domain**, **VLAN Pool**, and **VRF Lite IP Pool Range** information.
If you enable SDA connectivity for the site, you will need to configure additional settings as described in the SDA use case chapter of the [Cisco Multi-Site Configuration Guide for ACI Fabrics](#).
- k) (Optional) Enable **SR-MPLS Connectivity** for the site.
If the site is connected via an MPLS network, enable the **SR-MPLS Connectivity** knob and provide the Segment Routing global block (SRGB) range.

The Segment Routing Global Block (SRGB) is the range of label values reserved for Segment Routing (SR) in the Label Switching Database (LSD). These values are assigned as segment identifiers (SIDs) to SR-enabled nodes and have global significance throughout the domain.

The default range is 16000–23999.

If you enable MPLS connectivity for the site, you will need to configure additional settings as described in the "Sites Connected via SR-MPLS" chapter of the [Cisco Multi-Site Configuration Guide for ACI Fabrics](#).

Step 6 Configure inter-site connectivity between on-premises and cloud sites.

If you do not need to create inter-site connectivity between on-premises and cloud sites, for example if your deployment contains only cloud or only on-premises sites, skip this step.

When you configure underlay connectivity between on-premises and cloud sites, you need to provide an IPN device IP address to which the Cloud Network Controller's CSRs establish a tunnel and then configure the cloud site's infra settings.

- a) Click **+Add IPN Device** to specify an IPN device.
- b) From the dropdown, select one of the IPN devices you defined previously.

The IPN devices must be already defined in the **General Settings > IPN Devices** list, as described in [Configuring Infra: General Settings, on page 149](#)

- c) Configure inter-site connectivity for cloud sites.

Any previously configured connectivity from the cloud sites to this on-premises site will be displayed here, but any additional configuration must be done from the cloud site's side as described in [Configuring Infra for Cisco Cloud Network Controller Sites, on page 161](#).

What to do next

While you have configured all the required inter-site connectivity information, it has not been pushed to the sites yet. You need to deploy the configuration as described in [Deploying Infra Configuration, on page 167](#)

Configuring Infra: Pod Settings

This section describes how to configure Pod-specific settings in each site.

Procedure

-
- Step 1** Log in to the Cisco Nexus Dashboard Orchestrator GUI.
 - Step 2** In the left navigation menu, select **Infrastructure > Site Connectivity**.
 - Step 3** In the top right of the main pane, click **Configure**.
 - Step 4** In the left pane, under **Sites**, select a specific site.
 - Step 5** In the main window, select a Pod.
 - Step 6** In the right **Pod Properties** pane, add the Overlay Unicast TEP for the Pod.

This IP address is deployed on all spine switches that are part of the same Pod and used for sourcing and receiving VXLAN encapsulated traffic for Layer2 and Layer3 unicast communication.

Step 7 Click **+Add TEP Pool** to add an external routable TEP pool.

The external routable TEP pools are used to assign a set of IP addresses that are routable across the IPN to APIC nodes, spine switches, and border leaf nodes. This is required to enable Multi-Site architecture.

External TEP pools previously assigned to the fabric on APIC are automatically inherited by NDO and displayed in the GUI when the fabric is added to the Multi-Site domain.

Step 8 Repeat the procedure for every Pod in the site.

Configuring Infra: Spine Switches

This section describes how to configure spine switches in each site for Cisco Multi-Site. When you configure the spine switches, you are effectively establishing the underlay connectivity between the sites in your Multi-Site domain by configuring connectivity between the spines in each site and the ISN.

Prior to Release 3.5(1), underlay connectivity was establishing using OSPF protocol. In this release however, you can choose to use OSPF, BGP (IPv4 only), or a mixture of protocols, with some sites using OSPF and some using BGP for inter-site underlay connectivity. We recommend configuring either OSPF or BGP and not both, however if you configure both protocols, BGP will take precedence and OSPF will not be installed in the route table.

Procedure

Step 1 Log in to the Cisco Nexus Dashboard Orchestrator GUI.

Step 2 In the left navigation menu, select **Infrastructure > Site Connectivity**.

Step 3 In the top right of the main pane, click **Configure**.

Step 4 In the left pane, under **Sites**, select the specific on-premises site.

Step 5 In the main pane, select a spine switch within a pod.

Step 6 In the right **<Spine> Settings** pane, click **+Add Port**.

Step 7 In the **Add Port** window, provide the underlay connectivity information.

Any port already configured directly in APIC for IPN connectivity will be imported and shown in the list. For any new ports you want to configure from NDO, use the following the steps:

a) Provide general information:

- In the **Ethernet Port ID** field, enter the port ID, for example 1/29.

This is the interface which will be used to connect to the IPN.

- In the **IP Address** field, enter the IP address/netmask.

The Orchestrator creates a sub-interface with VLAN 4 with the specified IP ADDRESS under the specified PORT.

- In the **MTU** field, enter the MTU. You can specify either `inherit`, which would configure an MTU of 9150B, or choose a value between 576 and 9000.

MTU of the spine port should match MTU on IPN side.

Step 8 Choose the underlay protocol.

- a) Enable **OSPF** if you want to use OSPF protocol for underlay connectivity.

If you want to use BGP protocol for underlay connectivity instead, skip this part and provide the information required in the next substep.

- Set **OSPF** to *Enabled*.

The OSPF settings will become available.

- From the **OSPF Policy** dropdown, select the OSPF policy for the switch that you have configured in [Configuring Infra: On-Premises Site Settings, on page 154](#).

OSPF settings in the OSPF policy you choose should match on IPN side.

- For **OSPF Authentication**, you can pick either *none* or one of the following:

- MD5
- Simple

- Set **BGP** to *Disabled*.

- b) Enable **BGP** if you want to use BGP protocol for underlay connectivity.

If you're using OSPF protocol for underlay connectivity and have already configured it in the previous substep, skip this part.

Note

BGP IPv4 underlay is not supported in the following cases:

- If your Multi-Site domain contains one or more Cloud Network Controller sites, in which case you must use the OSPF protocol for intersite underlay connectivity for both on-prem to on-prem and on-prem to cloud sites.
- If you are using GOLF (Layer 3 EVPN services for fabric WAN) for WAN connectivity in any of your fabrics.

In the above cases, you must use OSPF in the Infra L3Out deployed on the spines.

- Set **OSPF** to *Disabled*.

We recommend configuring either OSPF or BGP and not both, however if you configure both protocols, BGP will take precedence and OSPF routes will not be installed in the route table because only EBGp adjacencies with the ISN devices are supported.

- Set **BGP** to *Enabled*.

The BGP settings will become available.

- In the **Peer IP** field, provide the IP address of this port's BGP neighbor.

Only IPv4 IP addresses are supported for BGP underlay connectivity.

- In the **Peer AS Number** field, provide the Autonomous System (AS) number of the BGP neighbor.

This release supports only EBGp adjacencies with the ISN devices.

- In the **BGP Password** field, provide the BGP peer password.

- Specify any additional options as required:

- `Bidirectional Forwarding Detection`—enables Bidirectional Forwarding Detection (BFD) protocol to detect faults on the physical link this port and the IPN device.
- `Admin State`—sets the admin state on the port to enabled.

Step 9 Repeat the procedure for every spine switch and port that connects to the IPN.



CHAPTER 14

Configuring Infra for Cisco Cloud Network Controller Sites

- [Refreshing Cloud Site Connectivity Information, on page 161](#)
- [Configuring Infra: Cloud Site Settings, on page 162](#)
- [Recovering from Cloud Network Controller Site Downtime , on page 164](#)

Refreshing Cloud Site Connectivity Information

Any infrastructure changes, such as CSR and Region addition or removal, require a Multi-Site fabric connectivity site refresh. This section describes how to pull up-to-date connectivity information directly from each site's APIC.

Procedure

-
- | | |
|---------------|---|
| Step 1 | Log in to the Cisco Nexus Dashboard Orchestrator GUI. |
| Step 2 | In the left navigation menu, select Infrastructure > Site Connectivity . |
| Step 3 | In the top right of the main pane, click Configure . |
| Step 4 | In the left pane, under Sites , select a specific site. |
| Step 5 | In the main window, click the Refresh button to discover any new or changed CSRs and regions. |
| Step 6 | Finally, click Yes to confirm and load the connectivity information.

This will discover any new or removed CSRs and regions. |
| Step 7 | Click Deploy to propagate the cloud site changes to other sites that have connectivity to it.

After you refresh a cloud site's connectivity and CSRs or regions are added or removed, you need to deploy infra configuration so other sites that have underlay connectivity to that cloud site get updated configuration. |
-

Configuring Infra: Cloud Site Settings

This section describes how to configure site-specific Infra settings for Cloud Network Controller sites.

Procedure

-
- Step 1** Log in to the Cisco Nexus Dashboard Orchestrator GUI.
- Step 2** In the left navigation menu, select **Infrastructure > Site Connectivity**.
- Step 3** In the top right of the main pane, click **Configure**.
- Step 4** In the left pane, under **Sites**, select a specific cloud site.
- Step 5** Provide the general **Inter-Site Connectivity** information.
- In the right **<Site> Settings** pane, select the **Inter-Site Connectivity** tab.
 - Enable the **Multi-Site** knob.
This defines whether the overlay connectivity is established between this site and other sites.
Note that the overlay configuration will not be pushed to sites which do not have the underlay intersite connectivity established as described in the next step.
 - (Optional) Specify the **BGP Password**.
- Step 6** Provide site-specific **Inter-Site Connectivity** information.
- In the right properties sidebar for the cloud site, click **Add Site**.
The **Add Site** window opens.
 - Under **Connected to Site**, click **Select a Site** and select the site (for example, `Site2`) to which you want to establish connectivity from the site you are configuring (for example, `Site1`).
Once you select the remote site, the **Add Site** window will update to reflect both directions of connectivity: **Site1 > Site2** and **Site2 > Site1**.
 - In the **Site1 > Site2** area, from the **Connection Type** dropdown, choose the type of connection between the sites.
The following options are available:
 - Public Internet**—connectivity between the two sites is established via the Internet.
This type is supported between any two cloud sites or between a cloud site and an on-premises site.
 - Private Connection**—connectivity is established using a private connection between the two sites.
This type is supported between a cloud site and an on-premises site.
 - Cloud Backbone**—connectivity is established using cloud backbone.
This type is supported between two cloud sites of the same type, such as Azure-to-Azure or AWS-to-AWS.
- If you have multiple types of sites (on-premises, AWS, and Azure), different pairs of site can use different connection type.
- Choose the **Protocol** that you want to use for connectivity between these two sites.

If using **BGP-EVPN** connectivity, you can optionally enable **IPSec** and choose which version of the Internet Key Exchange (IKE) protocol to use: IKEv1 (Version 1) or IKEv2 (Version 1) depending on your configuration.

- For `Public Internet` connectivity, IPsec is always enabled.
- For `Cloud Backbone` connectivity, IPsec is always disabled.
- For `Private Connection`, you can choose to enable or disable IPsec.

If using **BGP-IPv4** connectivity instead, you must provide an external VRF which will be used for route leaking configuration from the cloud site you are configuring.

After **Site1 > Site2** connectivity information is provided, the **Site2 > Site1** area will reflect the connectivity information in the opposite direction.

- e) Click **Save** to save the inter-site connectivity configuration.

When you save connectivity information from `Site1` to `Site2`, the reverse connectivity is automatically created from `Site2` to `Site1`, which you can see by selecting the other site and checking the **Inter-site Connectivity** information in the right sidebar.

- f) Repeat this step to add inter-site connectivity for other sites.

When you establish underlay connectivity from `Site1` to `Site2`, the reverse connectivity is done automatically for you.

However, if you also want to establish inter-site connectivity from `Site1` to `Site3`, you must repeat this step for that site as well.

Step 7 Provide **External Connectivity** information.

If you do not plan to configure connectivity to external sites or devices that are not managed by NDO, you can skip this step.

Detailed description of an external connectivity use case is available in the [Configuring External Connectivity from Cloud CSRs Using Nexus Dashboard Orchestrator](#) document.

- a) In the right **<Site> Settings** pane, select the **External Connectivity** tab.
b) Click **Add External Connection**.

The **Add External Connectivity** dialog will open.

- c) From the **VRF** dropdown, select the VRF you want to use for external connectivity.

This is the VRF which will be used to leak the cloud routes. The **Regions** section will display the cloud regions that contain the CSRs to which this configuration be applied.

- d) From the **Name** dropdown in the **External Devices** section, select the external device.

This is the external device you added in the **General Settings > External Devices** list during general infra configuration and must already be defined as described in [Configuring Infra: General Settings, on page 149](#).

- e) From the **Tunnel IKE Version** dropdown, pick the IKE version that will be used to establish the IPSec tunnel between the cloud site's CSRs and the external device.
f) (Optional) From the **Tunnel Subnet Pool** dropdown, choose one of the named subnet pools.

Named subnet pool are used to allocate IP addresses for IPSec tunnels between cloud site CSRs and external devices. If you do not provide any **named** subnet pools here, the **external** subnet pool will be used for IP allocation.

Providing a dedicated subnet pool for external device connectivity is useful for cases where a specific subnet is already being used to allocate IP addresses to the external router and you want to continue to use those subnets for IPSec tunnels for NDO and cloud sites.

If you want to provide a specific subnet pool for this connectivity, it must already be created as described in [Configuring Infra: General Settings, on page 149](#).

- g) (Optional) In the **Pre-Shared Key** field, provide the custom keys you want to use to establish the tunnel.
- h) If necessary, repeat the previous substeps for any additional external devices you want to add for the same external connection (same VRF).
- i) If necessary, repeat this step for any additional external connections (different VRFs).

Note that there's a one-to-one relationship for tunnel endpoints between CSRs and external devices, so while you can create additional external connectivity using different VRFs, you cannot create additional connectivity to the same external devices.

What to do next

While you have configured all the required inter-site connectivity information, it has not been pushed to the sites yet. You need to deploy the configuration as described in [Deploying Infra Configuration, on page 167](#)

Recovering from Cloud Network Controller Site Downtime

When Cloud Network Controller (formerly Cloud APIC) instance/VM goes down for any reason while still being managed by NDO, you may be unable to undeploy or delete any existing templates associated with that cloud site. In this case, attempting to forcefully unmanage the site in NDO can cause stale configuration and deployment errors even if the site recovers.

To recover from this:

Procedure

Step 1 Bring up the new Cloud Network Controller sites and re-register the cloud sites.

- a) Log in to ND.
- b) Open the admin console.
- c) Navigate to the **Sites** page.
- d) From the **Actions** menu next to the site you redeployed, choose **Edit Site**.
- e) Check the "Re-register site" checkbox.
- f) Provide the new site details.

You need to provide the new public IP address of site and login credentials.

- g) Click **Save** to re-register the site.

Once the connectivity status of the site shows **UP**, the site IPs in NDO are also updated and the new sites will be in 'managed' state.

Step 2 Undeploy the previously deployed templates for each schema.

- a) Log in to NDO.
- b) Navigate to **Application Management** and select **Schemas**.
- c) Click on a schema with the deployed templates.
- d) From the **Actions** menu next to the **Template Properties**, choose **Undeploy Template** and wait until the template is successfully undeployed.

Step 3 Refresh the site's infra configuration to ensure that the new Catalyst 8000V switches are added in NDO.

- a) Navigate to **Infrastructure** and select **Site Connectivity**.
- b) Click **Configure** at the top right of the screen.
- c) Select the cloud site under the **Sites** panel and click **Refresh**.
- d) Click **Deploy** on the top right of the screen and wait until all sites are successfully deployed.

Step 4 Redeploy all templates associated with this Cloud Network Controller site.

- a) Navigate to **Application Management** and select **Schemas**.
 - b) Click on a schema with the templates undeployed earlier.
 - c) Click **Deploy to Sites** and wait until the template is deployed.
-



CHAPTER 15

Deploying Infra Configuration for ACI Sites

- [Deploying Infra Configuration, on page 167](#)
- [Enabling Connectivity Between On-Premises and Cloud Sites, on page 168](#)

Deploying Infra Configuration

This section describes how to deploy the Infra configuration to each APIC site.

Procedure

Step 1

In the top right of the main pane, click **Deploy** and choose the appropriate option to deploy the configuration.

If you have configured only on-premises or only cloud sites, simply click **Deploy** to deploy the Infra configuration.

However, if you have both, on-premises and cloud site, the following additional options may be available:

- **Deploy & Download IPN Device Config files:** Pushes the configuration to both the on-premises APIC site and the Cloud Network Controller site and enables the end-to-end interconnect between the on-premises and the cloud sites.
In addition, this option downloads a zip file that contains configuration information that you will use to enable connectivity from the IPN devices to Cisco Cloud Services Router (CSR). A followup screen appears that allows you to select all or some of the configuration files to download.
- **Deploy & Download External Device Config files:** Pushes the configuration to both the Cloud Network Controller sites and enables the end-to-end interconnect between the cloud sites and external devices.
In addition, this option downloads a zip file that contains configuration information that you will use to enable connectivity from external devices to the Cisco Cloud Services Router (CSR) deployed in your cloud sites. A followup screen appears that allows you to select all or some of the configuration files to download.
- **Download IPN Device Config files only:** Downloads a zip file that contains configuration information that you will use to enable connectivity from the IPN devices to Cisco Cloud Services Router (CSR) without deploying the configuration.
- **Download External Device Config files only:** Downloads a zip file that contains configuration information that you will use to enable connectivity from external devices to Cisco Cloud Services Router (CSR) without deploying the configuration.

Step 2

In the confirmation window, click **Yes**.

The Deployment started, refer to left menu for individual site deployment status message will indicate that Infra configuration deployment began and you can verify each site's progress by the icon displayed next to the site's name in the left pane.

What to do next

The Infra overlay and underlay configuration settings are now deployed to all sites' controllers and cloud CSRs. The last remaining step is to configure your IPN devices with the tunnels for cloud CSRs as described in [Refreshing Site Connectivity Information, on page 153](#).

Enabling Connectivity Between On-Premises and Cloud Sites

If you have only on-premises or only cloud sites, you can skip this section.

This section describes how to enable connectivity between on-premises APIC sites and Cloud Network Controller sites.

By default, the Cisco Cloud Network Controller will deploy a pair of redundant Cisco Cloud Services Router 1000vs. The procedures in this section creates two tunnels, one IPsec tunnel from the on-premises IPsec device to each of these Cisco Cloud Services Router 1000vs. If you have multiple on-premises IPsec devices, you will need to configure the same tunnels to the CSRs on each of the on-premises devices.

The following information provides commands for Cisco Cloud Services Router 1000v as your on-premises IPsec termination device. Use similar commands if you are using a different device or platform.

Procedure

Step 1 Gather the necessary information that you will need to enable connectivity between the CSRs deployed in the cloud site and the on-premises IPsec termination device.

You can get the required configuration details using either the **Deploy & Download IPN Device config files** or the **Download IPN Device config files only** option in Nexus Dashboard Orchestrator as part of the procedures provided in [Deploying Infra Configuration, on page 167](#).

Step 2 Log into the on-premises IPsec device.

Step 3 Configure the tunnel for the *first* CSR.

Details for the first CSR are available in the configuration files for the ISN devices you downloaded from the Nexus Dashboard Orchestrator, but the following fields describe the important values for your specific deployment:

- *<first-csr-tunnel-ID>*—unique tunnel ID that you assign to this tunnel.
- *<first-csr-ip-address>*—public IP address of the third network interface of the first CSR.

The destination of the tunnel depends on the type of underlay connectivity:

- The destination of the tunnel is the public IP of the cloud router interface if the underlay is via public internet
- The destination of the tunnel is the private IP of the cloud router interface if the underlay is via private connectivity, such as DX on AWS or ER on Azure

- *<first-csr-preshared-key>*—preshared key of the first CSR.
- *<onprem-device-interface>*—interface that is used for connecting to the Cisco Cloud Services Router 1000v deployed in Amazon Web Services.
- *<onprem-device-ip-address>*—IP address for the *<interface>* interface that is used for connecting to the Cisco Cloud Services Router 1000v deployed in Amazon Web Services.
- *<peer-tunnel-for-onprem-IPsec-to-first-CSR>*—peer tunnel IP address for the on-premises IPsec device to the first cloud CSR.
- *<process-id>* —OSPF process ID.
- *<area-id>*—OSPF area ID.

The following example shows intersite connectivity configuration using the IKEv2 protocol supported starting with Nexus Dashboard Orchestrator, Release 3.3(1) and Cloud Network Controller, Release 5.2(1). If you are using IKEv1, the IPN configuration file you downloaded from NDO may look slightly differently, but the principle remains the same.

```
crypto ikev2 proposal ikev2-proposal-default
  encryption aes-cbc-256 aes-cbc-192 aes-cbc-128
  integrity sha512 sha384 sha256 sha1
  group 24 21 20 19 16 15 14 2
exit

crypto ikev2 policy ikev2-policy-default
  proposal ikev2-proposal-default
exit

crypto ikev2 keyring key-ikev2-infra:overlay-1-<first-csr-tunnel-id>
  peer peer-ikev2-keyring
    address <first-csr-ip-address>
    pre-shared-key <first-csr-preshared-key>
  exit
exit

crypto ikev2 profile ikev2-infra:overlay-1-<first-csr-tunnel-id>
  match address local interface <onprem-device-interface>
  match identity remote address <first-csr-ip-address> 255.255.255.255
  identity local address <onprem-device-ip-address>
  authentication remote pre-share
  authentication local pre-share
  keyring local key-ikev2-infra:overlay-1-<first-csr-tunnel-id>
  lifetime 3600
  dpd 10 5 on-demand
exit

crypto ipsec transform-set infra:overlay-1-<first-csr-tunnel-id> esp-gcm 256
  mode tunnel
exit

crypto ipsec profile infra:overlay-1-<first-csr-tunnel-id>
  set pfs group14
  set ikev2-profile ikev2-infra:overlay-1-<first-csr-tunnel-id>
  set transform-set infra:overlay-1-<first-csr-tunnel-id>
exit

interface tunnel 2001
  ip address <peer-tunnel-for-onprem-IPsec-to-first-CSR> 255.255.255.252
  ip virtual-reassembly
  tunnel source <onprem-device-interface>
  tunnel destination <first-csr-ip-address>
```

```

tunnel mode ipsec ipv4
tunnel protection ipsec profile infra:overlay-1-<first-csr-tunnel-id>
ip mtu 1400
ip tcp adjust-mss 1400
ip ospf <process-id> area <area-id>
no shut
exit

```

Example:

```

crypto ikev2 proposal ikev2-proposal-default
  encryption aes-cbc-256 aes-cbc-192 aes-cbc-128
  integrity sha512 sha384 sha256 sha1
  group 24 21 20 19 16 15 14 2
exit

crypto ikev2 policy ikev2-policy-default
  proposal ikev2-proposal-default
exit

crypto ikev2 keyring key-ikev2-infra:overlay-1-2001
  peer peer-ikev2-keyring
    address 52.12.232.0
    pre-shared-key 1449047253219022866513892194096727146110
  exit
exit

crypto ikev2 profile ikev2-infra:overlay-1-2001
  ! Please change GigabitEthernet1 to the appropriate interface
  match address local interface GigabitEthernet1
  match identity remote address 52.12.232.0 255.255.255.255
  identity local address 128.107.72.62
  authentication remote pre-share
  authentication local pre-share
  keyring local key-ikev2-infra:overlay-1-2001
  lifetime 3600
  dpd 10 5 on-demand
exit

crypto ipsec transform-set infra:overlay-1-2001 esp-gcm 256
  mode tunnel
exit

crypto ipsec profile infra:overlay-1-2001
  set pfs group14
  set ikev2-profile ikev2-infra:overlay-1-2001
  set transform-set infra:overlay-1-2001
exit

! These tunnel interfaces establish point-to-point connectivity between the on-prem device and the
cloud Routers
! The destination of the tunnel depends on the type of underlay connectivity:
! 1) The destination of the tunnel is the public IP of the cloud Router interface if the underlay is
via internet
! 2) The destination of the tunnel is the private IP of the cloud Router interface if the underlay
is via private
    connectivity like DX on AWS or ER on Azure

interface tunnel 2001
  ip address 5.5.1.26 255.255.255.252
  ip virtual-reassembly
  ! Please change GigabitEthernet1 to the appropriate interface
  tunnel source GigabitEthernet1
  tunnel destination 52.12.232.0
  tunnel mode ipsec ipv4

```



```
tunnel protection ipsec profile infra:overlay-1-2001
ip mtu 1400
ip tcp adjust-mss 1400
! Please update process ID according with your configuration
ip ospf 1 area 0.0.0.1
no shut
exit
```

Step 4 Repeat the previous step for the 2nd and any additional CSRs that you need to configure.

Step 5 Verify that the tunnels are up on your on-premises IPsec device.

Use the following command to display the status. If you do not see that both tunnels are shown as up, verify the information that you entered in the steps in this section to determine where you might have an issue. Do not proceed to the next section until you see that both tunnels are shown as up.

```
ISN_CSR# show ip interface brief | include Tunnel
```

Interface	IP-Address	OK?	Method	Status	Protocol
Tunnel1000	30.29.1.2	YES	manual	up	up
Tunnel1001	30.29.1.4	YES	manual	up	up



CHAPTER 16

CloudSec Encryption

- [Cisco ACI CloudSec Encryption, on page 173](#)
- [Requirements and Guidelines, on page 174](#)
- [CloudSec Encryption Terminology, on page 177](#)
- [CloudSec Encryption and Decryption Handling, on page 178](#)
- [CloudSec Encryption Key Allocation and Distribution, on page 180](#)
- [Configuring Cisco APIC for CloudSec Encryption, on page 182](#)
- [Enabling CloudSec Encryption in Nexus Dashboard Orchestrator, on page 185](#)
- [Verifying CloudSec Configuration on Switches, on page 186](#)
- [Rekey Process During Spine Switch Maintenance, on page 188](#)

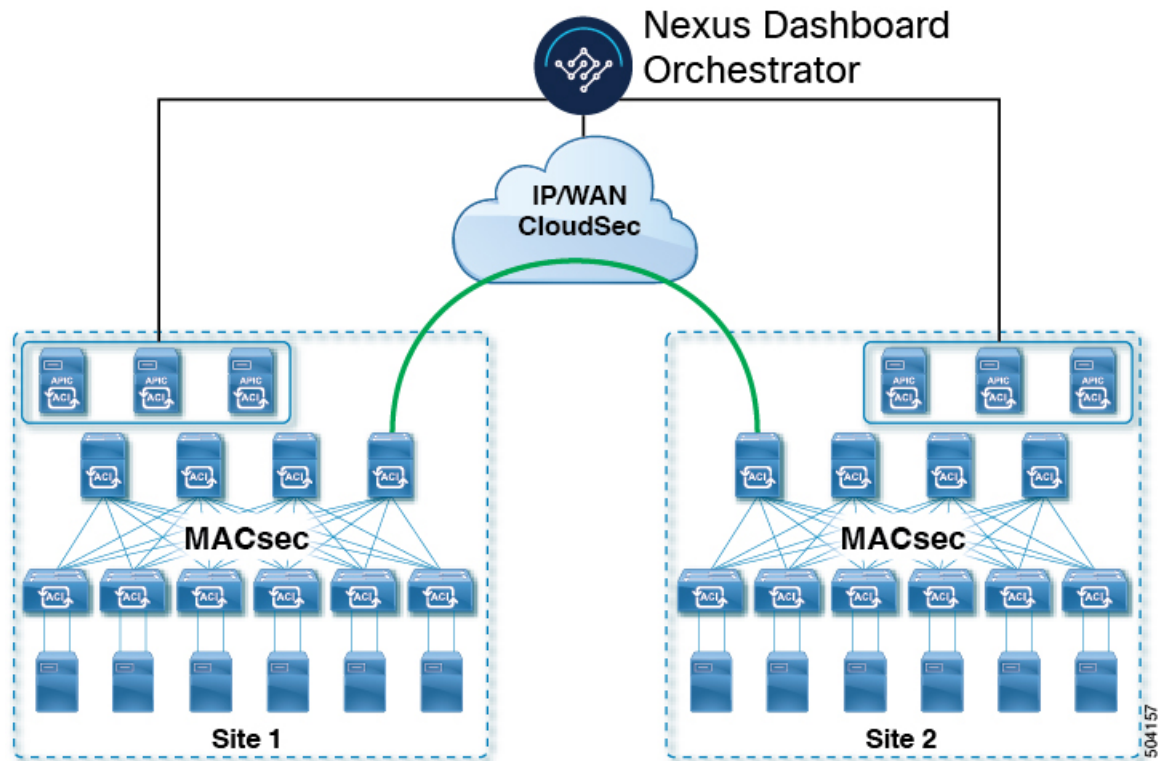
Cisco ACI CloudSec Encryption

As most Cisco ACI deployments are adopting the Multi-Site architecture to address disaster recovery and scale, the current security implementation using MACsec encryption within local site is becoming insufficient to guarantee data security and integrity across multiple sites connected by insecure external IP networks interconnecting separate fabrics. Nexus Dashboard Orchestrator Release 2.0(1) introduces the CloudSec Encryption feature designed to provide inter-site encryption of traffic.

Multi-Site topology uses three tunnel end-point (TEP) IP addresses (Overlay Multicast TEP, Overlay Unicast TEP, and External TEP Pool) to provide connectivity between sites. These TEP addresses are configured by the admin on Nexus Dashboard Orchestrator and pushed down to each site's Cisco APIC, which in turn configures them on the spine switches. These three addresses are used to determine when traffic is destined for a remote site, in which case an encrypted CloudSec tunnel is created between the two spine switches that provide physical connectivity between the two sites through the Inter-Site Network (ISN).

The following figure illustrates the overall encryption approach that combines MACsec for local site traffic and CloudSec for inter-site traffic encryption.

Figure 11: CloudSec Encryption



Requirements and Guidelines

When configuring CloudSec encryption, the following guidelines apply:

- CloudSec has been validated using a Nexus 9000 Inter-Site Network (ISN) infrastructure. If your ISN infrastructure is made up of different devices, or the devices are unknown (such as in the case of circuits purchased from a service provider), it is required that an ASR1K router is the first hop device directly connected to the ACI spine (with a separate pair of ASR1K devices deployed in each site), or the Nexus 9000 ISN network. The ASR1K router with padding-fixup enabled allows the CloudSec traffic to traverse any IP network between the sites.

To configure an ASR1K router:

1. Log in to the device.
2. Configure the UDP ports.

**Note**

- If you're running Release 3.7(1) or later and configure CloudSec to use the IANA-assigned port 8017, specify that port in the following command instead.
- You must have either an Advanced Enterprise or an Advanced IP Services license to run the **platform cloudsec padding-fixup** command on the ASR1K router.

```
ASR1K(config)# platform cloudsec padding-fixup dst-udp-port 9999
```

3. Verify the configuration.

In the following output, ensure that the port you configured in the previous step (8017 or 9999) is shown.

```
ASR1K# show platform software ip rp active cloudsec
CloudSec Debug: disabled
CloudSec UDP destination port: enabled
1st UDP destination port: 9999
2nd UDP destination port: 0
3rd UDP destination port: 0
```

```
ASR1K# show platform software ip fp active cloudsec
CloudSec Debug: disabled
CloudSec UDP destination port: enabled
1st UDP destination port: 9999
2nd UDP destination port: 0
3rd UDP destination port: 0
```

- If one or more spine switches are down when you attempt to disable CloudSec encryption, the disable process will not complete on those switches until the switches are up. This may result in packet drops on the switches when they come back up.

We recommend you ensure that all spine switches in the fabric are up or completely decommissioned before enabling or disabling CloudSec encryption.

- Beginning with Nexus Dashboard Orchestrator, Release 3.7(1), CloudSec encryption can be configured to use the IANA-assigned port.

By default, CloudSec uses a proprietary UDP port. Orchestrator releases 3.7(1) or later can be configured to use the official IANA-reserved port 8017 for CloudSec encryption between sites instead.

**Note**

The IANA-reserved port is supported for Cisco APIC sites running release 5.2(4) or later.

To change this setting, CloudSec must be disabled on all sites. If you want to enable IANA reserved port, but already have CloudSec encryption enabled for one or more of your sites, disable CloudSec for all sites, enable **IANA Reserve UDP Port** option, then re-enable CloudSec for the required sites.

- The CloudSec Encryption feature is not supported with the following features:
 - Precision Time Protocol (PTP)

- Remote Leaf Direct
- Virtual Pod (vPOD)
- SDA
- Remote Leaf or Multi-Pod configurations
- Intersite L3Out, if the sites are running Cisco APIC releases prior to 5.2(4).

CloudSec is supported with intersite L3Out for APIC sites running release 5.2(4) or later.

Requirements

The CloudSec encryption capability requires the following:

- Cisco ACI spine-leaf architecture with a Cisco APIC cluster for each site
- Cisco Nexus Dashboard Orchestrator to manage each site
- One **Advantage** or **Premier** license per each device (leaf only) in the fabric
- An add-on license **ACI-SEC-XF** per device for encryption if the device is a fixed spine
- An add-on license **ACI-SEC-XM** per device for encryption if the device is a modular spine

The following table provides the hardware platforms and the port ranges that are capable of CloudSec encryption.

Hardware Platform	Port Range
N9K-C9364C spine switches	Ports 49-64
N9K-C9332C spine switches	Ports 25-32
N9K-X9736C-FX line cards	Ports 29-36

If CloudSec is enabled for a site, but the encryption is not supported by the ports, a fault is raised with `unsupported-interface` error message.

CloudSec encryption's packet encapsulation is supported if Cisco QSFP-to-SFP Adapters (QSA), such as CVR-QSFP-SFP10G, is used with a supported optic. The full list of supported optics is available from the following link: <https://www.cisco.com/c/en/us/support/interfaces-modules/transceiver-modules/products-device-support-tables-list.html>.

Using IANA-Assigned Port and Orchestrator Downgrades

If you configure CloudSec encryption to use the IANA-assigned port as described in the following sections, there is a number of steps you have to take if you ever downgrade your Orchestrator service to a release prior to Release 3.7(1).

Before you downgrade your Nexus Dashboard Orchestrator to a release where IANA port is not supported:

1. Disable CloudSec encryption for all managed sites.
2. Disable the **IANA Reserved UDP Port** option in infra configuration settings.
3. Re-enable CloudSec encryption for all site where it was previously enabled.

4. Downgrade the Orchestrator services as you typically would.

CloudSec Encryption Terminology

CloudSec Encryption feature provides a secure upstream symmetric key allocation and distribution method for initial key and rekey requirements between sites. The following terminology is used in this chapter:

- **Upstream device** – The device that adds the CloudSec Encryption header and does the encryption of the VXLAN packet payload on transmission to a remote site using a locally generated symmetric cryptography key.
- **Downstream device** – The device that interprets the CloudSec Encryption header and does the decryption of the VXLAN packet payload on reception using the cryptography key generated by the remote site.
- **Upstream site** – The data center fabric that originates the encrypted VXLAN packets.
- **Downstream site** – The data center fabric that receives the encrypted packets and decrypts them.
- **TX Key** – The cryptography key used to encrypt the clear VXLAN packet payload. In ACI only one TX key can be active for all the remote sites.
- **RX Key** – The cryptography key used to decrypt the encrypted VXLAN packet payload. In ACI two RX keys can be active per remote site.

Two RX keys can be active at the same time because during the rekey process, the downstream sites will keep the old and the new RX keys after the new key deployment is finished for some duration to ensure that out of order packet deliveries with either key can be properly decrypted.

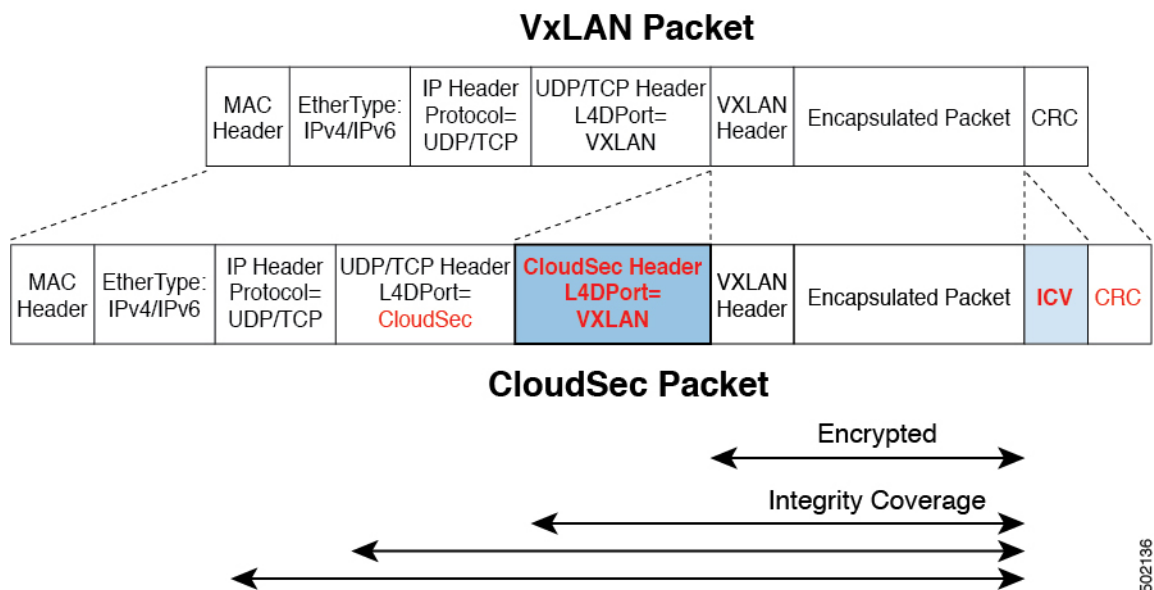
- **Symmetric Keys** – When the same cryptography key is used to encrypt (TX Key) and decrypt (RX Key) a packet stream by the upstream and downstream devices respectively.
 - **Rekey** – The process initiated by the upstream site to replace its old key with a newer key for all downstream sites after the old key expires.
 - **Secure Channel Identifier (SCI)** – A 64-bit identifier that represents a security association between the sites. It is transmitted in encrypted packet in CloudSec header and is used to derive the RX key on the downstream device for packet decryption.
 - **Association Number (AN)** – A 2-bit number (0, 1, 2, 3) that is sent in the CloudSec header of the encrypted packet and is used to derive the key at the downstream device in conjunction with the SCI for decryption. This allows multiple keys to be active at the downstream device to handle out of order packet arrivals with different keys from the same upstream device following a rekey operation.
- In ACI only two association number values (0 and 1) are used for the two active RX keys and only one association number value (0 or 1) is used for the TX key at any point in time.
- **Pre-shared key (PSK)** – One or more keys must be configured in the Cisco APIC GUI to be used as a random seed for generating the CloudSec TX and RX keys. If multiple PSK are configured, each rekey process will use the next PSK in order of their indexes; if no higher index PSK is available, a PSK with the lowest index will be used. Each PSK must be a hexadecimal string 64 characters long. Cisco APIC supports up to 256 pre-shared keys.

CloudSec Encryption and Decryption Handling

In order to provide a fully integrated, simple, and cost-effective solution that addresses both, data security and integrity, starting with Release 2.0(1), Multi-Site provides a CloudSec Encryption feature that allows for complete source-to-destination packet encryption between Multi-Site fabrics.

The following figure shows packet diagram before and after CloudSec encapsulation, followed by descriptions of the encryption and decryption processes:

Figure 12: CloudSec Packet



Packet Encryption

The following is a high level overview of how CloudSec handles outgoing traffic packets:

- The iVXLAN packets are filtered using outer IP header destination address field and Layer 4 destination port information and filtered packets are marked for encryption.
- The offset to use for encryption is calculated according to the fields of the packet. For example, the offset may vary based on whether there is a 802.1q VLAN or if the packet is an IPv4 or IPv6 packet.
The offset is automatically determined and is not visible to the user.
- The encryption keys are programmed in the hardware tables and are looked up from the table using the packet IP header.

Once the packet is marked for encryption, the encryption key is loaded, and the offset from the beginning of the packet where to start the encryption is known, the following additional steps are taken:

- The UDP destination port number is copied from the UDP header into a CloudSec field for recovery when the packet is decrypted.
- The UDP destination port number is overwritten to indicate that it is a CloudSec packet.

In releases prior to 3.7(1), the port is overwritten with a Cisco proprietary Layer-4 port number 9999.

In release 3.7(1) or later where you can configure CloudSec to use the IANA-assigned port 8017, the destination port number used is either 9999 or 8017 depending on whether you enabled this option.

- The UDP length field is updated to reflect the additional bytes that are being added.
- The CloudSec header is inserted directly after the UDP header.
- The Integrity Check Value (ICV) is inserted at the end of the packet, between the payload and the CRC.
- The ICV requires construction of a 128-bit initialization vector. For CloudSec, any use of the source MAC address for ICV purposes is replaced by a programmable value per SCI.
- CRC is updated to reflect the change in the contents of the packet.

Packet Decryption

The way CloudSec handles incoming packets is symmetric to the outgoing packets algorithm described above:

- If the received packet is a CloudSec packet, it is decrypted and the ICV is verified.

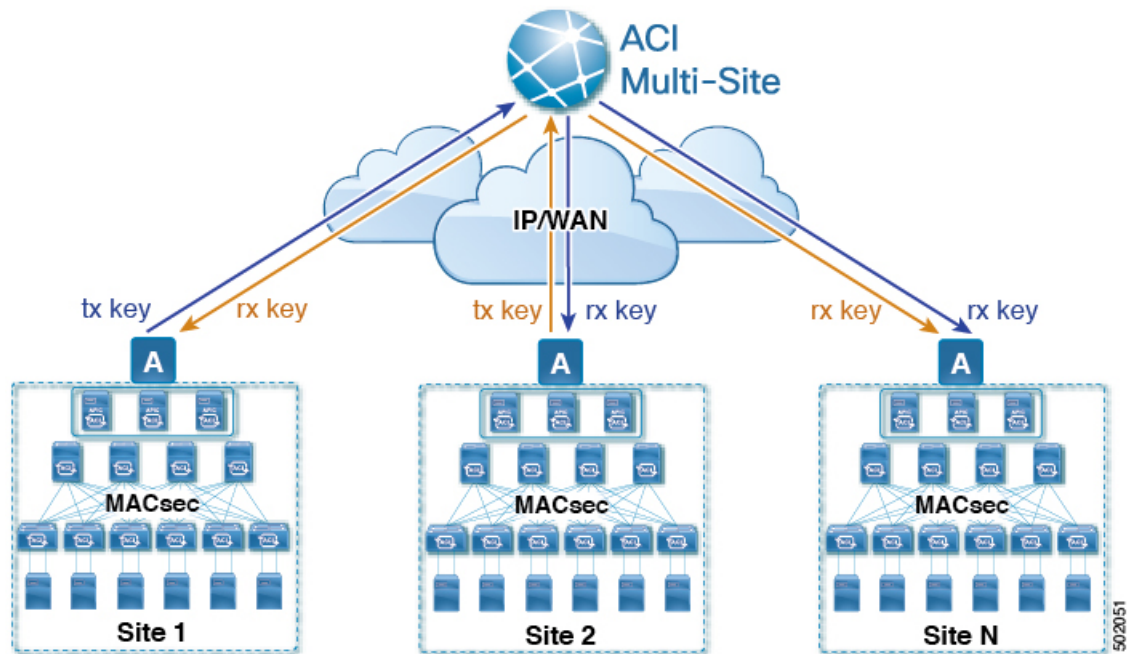
If ICV verification passed, the extra fields are removed, the UDP destination port number is moved from the CloudSec header to the UDP header, the CRC is updated, and the packet is forwarded to destination after decryption and CloudSec header removal. Otherwise the packet is dropped.

- The decryption key is retrieved from the key store using the received CloudSec packet outer IP header source address field, CloudSec header SCI, and AN number fields.
- If the packet is not a CloudSec packet, the packet is left unchanged.

CloudSec Encryption Key Allocation and Distribution

Initial Key Configuration

Figure 13: CloudSec Key Distribution



The following is a high level overview of the CloudSec encryption key initial allocation and distribution process illustrated by the figure above:

- The upstream site's Cisco APIC generates a local symmetric key intended to be used for data encryption of VXLAN packets transmitted from its site. The same key that is used by the upstream site for encryption is used for decryption of the packets on the downstream remote receiving sites.

Every site is an upstream site for the traffic it transmits to other sites. If multiple sites exist, each site generates its own site-to-site key and use that key for encryption before transmitting to the remote site.

- The generated symmetric key is pushed to the Nexus Dashboard Orchestrator (NDO) by the upstream site's Cisco APIC for distribution to downstream remote sites.
- The NDO acts as a message broker and collects the generated symmetric key from the upstream site's Cisco APIC, then distributes it to downstream remote sites' Cisco APICs.

The keys are encrypted using Key Encryption Key (KEK) and are distributed via a TLS-based channel.

- Each downstream site's Cisco APIC configures the received key as RX key on the local spine switches which are intended to receive the traffic from the upstream site that generated the key.
- Each downstream site's Cisco APIC also collects the deployment status of the RX Key from the local spine switches and then pushes it to the NDO.
- The NDO relays the key deployment status from all downstream remote sites back to the upstream site's Cisco APIC.

- The upstream site's Cisco APIC checks if the key deployment status received from all downstream remote sites is successful.
 - If the deployment status received from a downstream device is successful, the upstream site deploys the local symmetric key as its TX key on the spine switches to enable encryption of the VXLAN packets that are sent to the downstream site.
 - If the deployment status received from a downstream device is failed, a fault is raised on the Cisco APIC site where it failed and it is handled based on the "secure mode" setting configured on the NDO. In "must secure" mode the packets are dropped and in the "should secure" mode the packets are sent clear (unencrypted) to the destination site.



Note In current release, the mode is always set to "should secure" and cannot be changed.

Rekey Process

Each generated TX/RX key expires after a set amount of time, by default key expiry time is set to 15 minutes. When the initial set of TX/RX keys expires, a rekey process takes place.

The same general key allocation and distribution flow applies for the rekey process. The rekey process follows the "make before break" rule, in other words all the RX keys on the downstream sites are deployed before the new TX key is deployed on the upstream site. To achieve that, the upstream site will wait for the new RX key deployment status from the downstream sites before it configures the new TX key on the local upstream site's devices.

If any downstream site reports a failure status in deploying the new RX key, the rekey process will be terminated and the old key will remain active. The downstream sites will also keep the old and the new RX keys after the new key deployment is finished for some duration to ensure that out of order packet deliveries with either key can be properly decrypted.



Note Special precautions must be taken in regards to rekey process during spine switch maintenance, see [Rekey Process During Spine Switch Maintenance, on page 188](#) for details.

Rekey Process Failure

In case of any downstream site failing to deploy the new encryption key generated by the rekey process, the new key is discarded and the upstream device will continue to use the previous valid key as TX key. This approach keeps the upstream sites from having to maintain multiple TX keys per set of downstream sites. However, this approach may also result in the rekey process being delayed if the rekey deployment failures continue to occur with any one of the downstream sites. It is expected that the Multi-Site administrator will take action to fix the issue of the key deployment failure for the rekey to succeed.

Cisco APIC's Role in Key Management

The Cisco APIC is responsible for key allocation (both, initial key and rekey distribution), collection of the key deployment status messages from the spine switches, and notification of the Nexus Dashboard Orchestrator about each key's status for distribution to other sites.

Nexus Dashboard Orchestrator's Role in Key Management

The Nexus Dashboard Orchestrator is responsible for collecting the TX keys (both, initial key and subsequent rekeys) from the upstream site and distributing it to all downstream sites for deployment as RX keys. The NDO also collects the RX key deployment status information from the downstream sites and notifies the upstream site in order for it to update the TX key on successful RX key deployment status.

Upstream Model

In contrast to other technologies, such as MPLS, that use downstream key allocation, CloudSec's upstream model provides the following advantages:

- The model is simple and operationally easier to deploy in the networks.
- The model is preferred for Multi-Site use cases.
- It provides advantages for multicast traffic as it can use the same key and CloudSec header for each copy of the replicated packet transmitted to multiple destination sites. In downstream model each copy would have to use a different security key for each site during encryption.
- It provides easier troubleshooting in case of failures and better traceability of packets from the source to destination consistently for both, unicast and multicast replicated packets.

Configuring Cisco APIC for CloudSec Encryption

You must configure one or more Pre-Shared Keys (PSK) to be used by the Cisco APIC for generating the CloudSec encryption and decryption keys. The PSK are used as a random seed during the re-key process. If multiple PSK are configured, each re-key process will use the next PSK in order of their indexes; if no higher index PSK is available, a PSK with the lowest index will be used.

Because PSK is used as a seed for encryption key generation, configuring multiple PSK provides additional security by lowering the over-time vulnerability of the generated encryption keys.



Note If no pre-shared key is configured on the Cisco APIC, CloudSec will not be enabled for that site. In that case, turning on CloudSec setting in Multi-Site will raise a fault.

If at any time you wish to refresh a previously added PSK with a new one, simply repeat the procedure as if you were adding a new key, but specify an existing index.

You can configure one or more pre-shared keys in one of three ways:

- Using the Cisco APIC GUI, as described in [Configuring Cisco APIC for CloudSec Encryption Using GUI, on page 183](#)
- Using the Cisco APIC NX-OS Style CLI, as described in [Configuring Cisco APIC for CloudSec Encryption Using NX-OS Style CLI, on page 183](#)
- Using the Cisco APIC REST API, as described in [Configuring Cisco APIC for CloudSec Encryption Using REST API, on page 184](#)

Configuring Cisco APIC for CloudSec Encryption Using GUI

This section describes how to configure one or more pre-shared keys (PSK) using the Cisco APIC GUI.

Procedure

Step 1 Log in to APIC.

Step 2 Navigate to **Tenants > infra > Policies > CloudSec Encryption**

Step 3 Specify the **SA Key Expiry Time**.

This option specifies how long each key is valid (in minutes). Each generated TX/RX key expires after the specified amount of time triggering a re-key process. The expiration time can be between 5 and 1440 minutes.

Step 4 Click the + icon in the **Pre-Shared Keys** table.

Step 5 Specify the **Index** of the pre-shared key you are adding and then the **Pre-Shared Key** itself.

The **Index** field specifies the order in which the pre-shared keys are used. After the last (highest index) key is used, the process will continue with the first (lowest index) key. Cisco APIC supports up to 256 pre-shared keys, so the PSK index value must be between 1 and 256.

Each **Pre-Shared Key** must be a hexadecimal string 64 characters long.

Configuring Cisco APIC for CloudSec Encryption Using NX-OS Style CLI

This section describes how to configure one or more pre-shared keys (PSK) using the Cisco APIC NX-OS Style CLI.

Procedure

Step 1 Log in to the Cisco APIC NX-OS style CLI.

Step 2 Enter configuration mode.

Example:

```
apicl# configure
apicl (config)#
```

Step 3 Enter configuration mode for the default CloudSec profile.

Example:

```
apicl(config)# template cloudsec default
apicl(config-cloudsec)#
```

Step 4 Specify the Pre-Shared Keys (PSK) expiration time.

This option specifies how long each key is valid (in minutes). Each generated TX/RX key expires after the specified amount of time triggering a re-key process. The expiration time can be between 5 and 1440 minutes.

Example:

```
apicl(config-cloudsec)# sakexpirytime <duration>
```



```
</cloudsecIfPol>  
</fvTenant>
```

Enabling CloudSec Encryption in Nexus Dashboard Orchestrator

The CloudSec encryption can be enabled or disabled for each site individually. However, the communications between two sites will be encrypted only if the feature is enabled on both sites.

Before you begin

Before you enable the CloudSec encryption between two or more sites, you must have completed the following tasks:

- Installed and configured the Cisco APIC clusters in multiple sites, as described in *Cisco APIC Installation, Upgrade, and Downgrade Guide*
- Installed and configured Nexus Dashboard Orchestrator, as described in *Cisco Nexus Dashboard Orchestrator Installation and Upgrade Guide*.
- Added each Cisco APIC site to the Nexus Dashboard Orchestrator, as described in *Cisco Multi-Site Configuration Guide*.

Procedure

- Step 1** Log in to the Nexus Dashboard Orchestrator.
- Step 2** From the left-hand navigation menu, select the **Infrastructure > Site Connectivity**.
- Step 3** Click the **Configure** button in the top right of the main pane.
- Step 4** (Optional) In the **General Settings** page's **Control Plane Configuration** tab, enable the **IANA Reserved UDP Port** option.

By default, CloudSec uses a proprietary UDP port. This option allows you to configure CloudSec to use the official IANA-reserved port 8017 for CloudSec encryption between sites.

Note

The IANA-reserved port is supported for Cisco APIC sites running release 5.2(4) or later.

To change this setting, CloudSec must be disabled on all sites. If you want to enable IANA reserved port, but already have CloudSec encryption enabled for one or more of your sites, disable CloudSec for all sites, enable **IANA Reserve UDP Port** option, then re-enable CloudSec for the required sites.

- Step 5** From the left-hand sidebar, select the site for which you want to change the CloudSec configuration.
- Step 6** In the right-hand sidebar, toggle the **CloudSec Encryption** setting to enable or disable the CloudSec encryption feature for the site.
-

Verifying CloudSec Configuration on Switches

The following command allows you to see the current CloudSec configuration that was deployed to the spine switch after you enable CloudSec encryption from your Nexus Dashboard Orchestrator.

Procedure

Step 1 Log in to your spine switch.

Step 2 Run the `show cloudsec sa interface all` command to show CloudSec configuration.

In the following output, ensure that for each `Interface`:

- The `Operational Status` value shows `UP`.
- The `Control` value is the same across all interfaces in all CloudSec-enabled sites as it indicates the UDP port currently in use for CloudSec encryption.

The following example shows the default Cisco-proprietary UDP port (`deprecatedUdpPort`). If you configure CloudSec to use the IANA-assigned port 8017, the **Control** field will display `ianaUdpPort` instead.

```
spinel# show cloudsec sa interface all
=====
Interface: Eth1/49.49(0x1a030031) Physical Interface: Eth1/49(0x1a030000)
  Operational Status: UP Retry: Off Control: deprecatedUdpPort
-----
Site-Id: 2 Peer: 200.200.204.0/24 Type: ext-routable-tep-pool Operational Status: UP
Pod-Id: 1
-----
TX Key: ***** Assoc Num: 1 Sci: 0x10002
Cipher Suite: gcm-aes-xpn-256 Rekey Num: 6826 Oper Rekey Num: 6826
Hardware Index: 0 Operational Status: UP Control: NONE
Last Updated: PST 2022-01-11 23:26:37.520-08:00 Retry: Off
Uptime: 11 hours 30 mins 45 secs
-----
Site-Id: 2 Peer: 200.200.202.1/32 Type: msite-unicast-tep Operational Status: UP
-----
TX Key: ***** Assoc Num: 1 Sci: 0x10002
Cipher Suite: gcm-aes-xpn-256 Rekey Num: 6826 Oper Rekey Num: 6826
Hardware Index: 2 Operational Status: UP Control: NONE
Last Updated: PST 2022-01-11 23:26:37.563-08:00 Retry: Off
Uptime: 11 hours 30 mins 45 secs
RX Key: ***** Assoc Num: 1 Sci: 0x20001
Cipher Suite: gcm-aes-xpn-256 Rekey Num: 6826 Oper Rekey Num: 6826
Hardware Index: 3 Operational Status: UP Control: NONE
Last Updated: PST 2022-01-11 23:26:37.442-08:00 Retry: Off
Uptime: 11 hours 30 mins 45 secs
RX Key: ***** Assoc Num: 0 Sci: 0x20001
Cipher Suite: gcm-aes-xpn-256 Rekey Num: 6827 Oper Rekey Num: 6827
Hardware Index: 2 Operational Status: UP Control: NONE
Last Updated: PST 2022-01-11 23:26:37.453-08:00 Retry: Off
Uptime: 11 hours 30 mins 45 secs
-----
Site-Id: 2 Peer: 200.200.201.1/32 Type: msite-multicast-tep Operational Status: UP
-----
TX Key: ***** Assoc Num: 1 Sci: 0x10002
Cipher Suite: gcm-aes-xpn-256 Rekey Num: 6826 Oper Rekey Num: 6826
```



```

Hardware Index: 1 Operational Status: UP Control: NONE
Last Updated: PST 2022-01-11 23:26:37.549-08:00 Retry: Off
Uptime: 11 hours 30 mins 45 secs
RX Key: ***** Assoc Num: 1 Sci: 0x20001
Cipher Suite: gcm-aes-xpn-256 Rekey Num: 6826 Oper Rekey Num: 6826
Hardware Index: 1 Operational Status: UP Control: NONE
Last Updated: PST 2022-01-11 23:26:36.501-08:00 Retry: Off
Uptime: 11 hours 30 mins 46 secs
RX Key: ***** Assoc Num: 0 Sci: 0x20001
Cipher Suite: gcm-aes-xpn-256 Rekey Num: 6827 Oper Rekey Num: 6827
Hardware Index: 0 Operational Status: UP Control: NONE
Last Updated: PST 2022-01-11 23:26:37.495-08:00 Retry: Off
Uptime: 11 hours 30 mins 45 secs

=====
Interface: Eth1/50.50(0x1a031032) Physical Interface: Eth1/50(0x1a031000)
Operational Status: UP Retry: Off Control: deprecatedUdpPort
-----
Site-Id: 2 Peer: 200.200.204.0/24 Type: ext-routable-tep-pool Operational Status: UP
Pod-Id: 1
-----
TX Key: ***** Assoc Num: 1 Sci: 0x10002
Cipher Suite: gcm-aes-xpn-256 Rekey Num: 6826 Oper Rekey Num: 6826
Hardware Index: 1 Operational Status: UP Control: NONE
Last Updated: PST 2022-01-11 23:26:37.577-08:00 Retry: Off
Uptime: 11 hours 30 mins 45 secs
-----
Site-Id: 2 Peer: 200.200.201.1/32 Type: msite-multicast-tep Operational Status: UP
-----
TX Key: ***** Assoc Num: 1 Sci: 0x10002
Cipher Suite: gcm-aes-xpn-256 Rekey Num: 6826 Oper Rekey Num: 6826
Hardware Index: 0 Operational Status: UP Control: NONE
Last Updated: PST 2022-01-11 23:26:37.537-08:00 Retry: Off
Uptime: 11 hours 30 mins 45 secs
RX Key: ***** Assoc Num: 1 Sci: 0x20001
Cipher Suite: gcm-aes-xpn-256 Rekey Num: 6826 Oper Rekey Num: 6826
Hardware Index: 1 Operational Status: UP Control: NONE
Last Updated: PST 2022-01-11 23:26:36.463-08:00 Retry: Off
Uptime: 11 hours 30 mins 46 secs
RX Key: ***** Assoc Num: 0 Sci: 0x20001
Cipher Suite: gcm-aes-xpn-256 Rekey Num: 6827 Oper Rekey Num: 6827
Hardware Index: 0 Operational Status: UP Control: NONE
Last Updated: PST 2022-01-11 23:26:37.416-08:00 Retry: Off
Uptime: 11 hours 30 mins 45 secs

-----
Site-Id: 2 Peer: 200.200.202.1/32 Type: msite-unicast-tep Operational Status: UP
-----
TX Key: ***** Assoc Num: 1 Sci: 0x10002
Cipher Suite: gcm-aes-xpn-256 Rekey Num: 6826 Oper Rekey Num: 6826
Hardware Index: 2 Operational Status: UP Control: NONE
Last Updated: PST 2022-01-11 23:26:37.593-08:00 Retry: Off
Uptime: 11 hours 30 mins 45 secs
RX Key: ***** Assoc Num: 0 Sci: 0x20001
Cipher Suite: gcm-aes-xpn-256 Rekey Num: 6827 Oper Rekey Num: 6827
Hardware Index: 2 Operational Status: UP Control: NONE
Last Updated: PST 2022-01-11 23:26:37.481-08:00 Retry: Off
Uptime: 11 hours 30 mins 45 secs
RX Key: ***** Assoc Num: 1 Sci: 0x20001
Cipher Suite: gcm-aes-xpn-256 Rekey Num: 6826 Oper Rekey Num: 6826
Hardware Index: 3 Operational Status: UP Control: NONE

```

Last Updated: PST 2022-01-11 23:26:37.507-08:00 Retry: Off
Uptime: 11 hours 30 mins 45 secs

Rekey Process During Spine Switch Maintenance

The following is a summary of the CloudSec rekey process during typical maintenance scenarios for the spine switches where the feature is enabled:

- **Normal Decommissioning** – CloudSec rekey process stops automatically whenever a CloudSec-enabled spine switch is decommissioned. Rekey process will not start again until the decommissioned node is commissioned back or the decommissioned node ID is removed from the Cisco APIC
- **Spine Switch Software Upgrade** – CloudSec rekey process stops automatically if a spine switch is reloaded due to software upgrade. Rekey process will resume after the spine switch comes out of reload.
- **Maintenance (GIR mode)** – CloudSec rekey process must be manually stopped using the instructions provided in [Disabling and Re-Enabling Re-Key Process Using NX-OS Style CLI, on page 188](#). Rekey can be enabled back only after the node is ready to forward traffic again.
- **Decommissioning and Removal from Cisco APIC** – CloudSec rekey process must be manually stopped using the instructions provided in [Disabling and Re-Enabling Re-Key Process Using NX-OS Style CLI, on page 188](#). Rekey can be enabled back only after the node is removed from Cisco APIC.

Disabling and Re-Enabling Re-Key Process Using NX-OS Style CLI

It is possible to manually stop and restart the re-key process. You may be required to manually control the re-key process in certain situations, such as switch decommissioning and maintenance. This section describes how to toggle the setting using Cisco APIC NX-OS Style CLI.

Procedure

Step 1 Log in to the Cisco APIC NX-OS style CLI.

Step 2 Enter configuration mode.

Example:

```
apic1# configure
apic1(config)#
```

Step 3 Enter configuration mode for the default CloudSec profile.

Example:

```
apic1(config)# template cloudsec default
apic1(config-cloudsec)#
```

Step 4 Stop or restart the re-key process.

To stop the re-key process:

Example:

```
apic1(config-cloudsec)# stoprekey yes
```

To restart the re-key process:

Example:

```
apic1(config-cloudsec)# stoprekey no
```

Disabling and Re-Enabling Re-Key Process Using REST API

It is possible to manually stop and restart the re-key process. You may be required to manually control the re-key process in certain situations, such as switch decommissioning and maintenance. This section describes how to toggle the setting using Cisco APIC REST API.

Procedure

Step 1 You can disable the rekey process using the following XML message.

Example:

```
<fvTenant annotation="" descr="" dn="uni/tn-infra" name="infra" nameAlias="" ownerKey="" ownerTag="">
  <cloudsecIfPol descr="cloudsecifp" name="default" sakExpiryTime="10" stopRekey= "true" status=""
/>
</fvTenant>
```

Step 2 You can enable the rekey process using the following XML message.

Example:

```
<fvTenant annotation="" descr="" dn="uni/tn-infra" name="infra" nameAlias="" ownerKey="" ownerTag="">
  <cloudsecIfPol descr="cloudsecifp" name="default" sakExpiryTime="10" stopRekey= "false" status=""
/>
</fvTenant>
```



PART IV

Features and Use Cases

- [DHCP Relay, on page 193](#)
- [EPG Preferred Group, on page 201](#)
- [Intersite L3Out, on page 205](#)
- [Intersite L3Out with PBR, on page 225](#)
- [Layer 3 Multicast, on page 251](#)
- [QoS Preservation Across IPN, on page 261](#)
- [SD-Access and ACI Integration, on page 267](#)
- [SD-WAN Integration, on page 285](#)
- [Multi-Site and SR-MPLS L3Out Handoff, on page 293](#)
- [vzAny Contracts, on page 313](#)



DHCP Relay

- [DHCP Relay Policy, on page 193](#)
- [Guidelines and Limitations, on page 193](#)
- [Creating DHCP Relay Policies, on page 194](#)
- [Creating DHCP Option Policies, on page 196](#)
- [Assigning DHCP Policies, on page 197](#)
- [Creating DHCP Relay Contract, on page 198](#)
- [Verifying DHCP Relay Policies in APIC, on page 199](#)
- [Editing or Deleting Existing DHCP Policies, on page 200](#)

DHCP Relay Policy

Typically, when your DHCP server is located under an EPG, all the endpoints in that EPG have access to it and can obtain the IP addresses via DHCP. However, in many deployment scenarios, the DHCP server may not exist in the same EPG, BD, or VRF as all the clients that require it. In these cases a DHCP relay can be configured to allow endpoints in one EPG to obtain IP addresses via DHCP from a server that is located in another EPG/BD deployed in a different site or even connected externally to the fabric and reachable via an L3Out connection.

You can create the DHCP `Relay` policy in the Orchestrator GUI to configure the relay. Additionally, you can choose to create a DHCP `Option` policy to configure additional options you can use with the relay policy to provide specific configuration details. For all available DHCP options refer to [RFC 2132](#).

When creating a DHCP relay policy, you specify an EPG (for example, `epg1`) or external EPG (for example, `ext-epg1`) where the DHCP server resides. After you create the DHCP policy, you associate it with a bridge domain, which in turn is associated with another EPG (for example, `epg2`) allowing the endpoints in that EPG to reach the DHCP server. Finally, you create a contract between the relay EPG (`epg1` or `ext-epg1`) and application EPG (`epg2`) to allow communication. The DHCP policies you create are pushed to the APIC when the bridge domain to which the policy is associated is deployed to a site.

Guidelines and Limitations

The DHCP relay policies are supported with the following caveats:

- DHCP relay policies are supported for fabrics running Cisco APIC Release 4.2(1) or later.
- The DHCP servers must support DHCP Relay Agent Information Option (Option 82).

When an ACI fabric acts as a DHCP relay, it inserts the DHCP Relay Agent Information Option in DHCP requests that it proxies on behalf of clients. If a response (DHCP offer) comes back from a DHCP server without Option 82, it is silently dropped by the fabric.

- DHCP relay policies are supported in user tenants or the `common` tenant only. DHCP policies are not supported for the `infra` or `mgmt` tenants.

When configuring shared resources and services in the ACI fabric, we recommend creating those resources in the `common` tenant, that way they can be used by any user tenant.

- DHCP relay server must be in the same user tenant as the DHCP clients or in the `common` tenant.

The server and the clients cannot be in different user tenants.

- DHCP relay policies can be configured for the primary SVI interface only.

If the bridge domain to which you assign a relay policy contains multiple subnets, the first subnet you add becomes the primary IP address on the SVI interface, while additional subnets are configured as secondary IP addresses. In certain scenarios, such as importing a configuration with a bridge domain with multiple subnets, the primary address on the SVI may change to one of the secondary addresses, which would break the DHCP relay for that bridge domain.

You can use the `show ip interface vrf all` command to verify IP address assignments for the SVI interfaces.

- If you make changes to the DHCP policy after you have assigned it to a bridge domain and deployed the bridge domain to one or more sites, you will need to re-deploy the bridge domain for the DHCP policy changes to be updated on each site's APIC.
- For inter-VRF DHCP relay with the DHCP server reachable via an L3Out, DHCP relay packets must use site-local L3Out to reach the DHCP server. Packets using an L3Out in a different site (Intersite L3Out) to reach the DHCP server is not supported.
- The following DHCP relay configurations are not supported:
 - DHCP relay label on L3Out interfaces.
 - Importing existing DHCP policies from APIC.
 - DHCP relay policy configuration in Global Fabric Access Policies is not supported
 - Multiple DHCP servers within the same DHCP relay policy and EPG.

If you configure multiple providers under the same DHCP relay policy, they must be in different EPGs or external EPGs.

Creating DHCP Relay Policies

This section describes how to create a DHCP relay policy.



Note

If you make changes to the DHCP policy after you have assigned it to a bridge domain and deployed the bridge domain to one or more sites, you will need to re-deploy the bridge domain for the DHCP policy changes to update on each site's APIC.

Before you begin

You must have the following:

- A DHCP server set up and configured in your environment.
- If the DHCP server is part of an application EPG, that EPG must be already created in the Nexus Dashboard Orchestrator.
- If the DHCP server is external to the fabric, the external EPG associated to the L3Out that is used to access the DHCP server must be already created.

Procedure

Step 1 Log in to your Nexus Dashboard and open the Nexus Dashboard Orchestrator service.

Step 2 Create a new Tenant Policy.

- a) From the left navigation pane, choose **Application Management > Tenant Policies**.
- b) On the **Tenant Policy Templates** page, click **Add Tenant Policy Template**.
- c) In the Tenant Policies page's right properties sidebar, provide the **Name** for the tenant.
- d) From the **Select a Tenant** dropdown, choose the tenant with which you want to associate this template.

All the policies you create in this template as described in the following steps will be associated with the selected tenant and deployed to it when you push the template to a specific site.

Step 3 Create a DHCP Relay Policy.

- a) From the **+Create Object** dropdown, select **DHCP Relay Policy**.
- b) In the right properties sidebar, provide the **Name** for the policy.
- c) (Optional) Click **Add Description** and provide a description for the policy.
- d) Click **Add Provider** to configure the DHCP server to which you want to relay the DHCP requests originated by the endpoints.
- e) Select the provider type.

When adding a relay policy, you can choose one of the following two types:

- **Application EPG**—specifies the application EPG that includes the DHCP server to which you want to relay the DHCP requests.
- **L3 External Network**—specifies the External EPG associated to the L3Out that is used to access the network external to the fabric where the DHCP server is connected.

Note

You can select any EPG or external EPG that has been created in the Orchestrator and assigned to the tenant you specified, even if you have not yet deployed it to sites. If you select an EPG that hasn't been deployed, you can still complete the DHCP relay configuration, but you will need to deploy the EPG before the relay is available for use.

- f) Click **Select an Application EPG** or **Select an External EPG** (based on the provider type you selected) and choose the provider EPG.
- g) In the **DHCP Server Address** field, provide the IP address of the DHCP server.
- h) Enable the **DHCP Server VRF Preference** option if required.

This feature was introduced in Cisco APIC release 5.2(4). For more information on the use cases where it is required refer to the [Cisco APIC Basic Configuration Guide](#).

- i) Click **OK** to save the provider information.
 - j) Repeat the previous substeps for any additional providers in the same DHCP Relay policy.
 - k) Repeat this step to create any additional DHCP Relay policies.
-

Creating DHCP Option Policies

This section describes how to create a DHCP option policy. DHCP options are appended to the end of the messages that DHCP servers and clients exchange and can be used to provide additional configuration information to your DHCP server. Each DHCP option has a specific code that you must provide when adding the option policy. For a complete list of DHCP options and codes, see [RFC 2132](#).

Before you begin

You must have the following already configured:

- A DHCP server set up and configured in your environment.
- An EPG that contains the DHCP server already created in the Nexus Dashboard Orchestrator.
- A DHCP Relay policy created, as described in [Creating DHCP Relay Policies, on page 194](#).

Procedure

Step 1 Log in to your Nexus Dashboard and open the Nexus Dashboard Orchestrator service.

Step 2 Create a new or update an existing Tenant Policy.

- a) From the left navigation pane, choose **Application Management > Tenant Policies**.
- b) On the **Tenant Policy Templates** page, select an existing policy or click **Add Tenant Policy Template**.
- c) If creating a new policy, in the Tenant Policies page's right properties sidebar, provide the **Name** for the tenant.
- d) If creating a new policy, from the **Select a Tenant** dropdown, choose the tenant with which you want to associate this template.

All the policies you create in this template as described in the following steps will be associated with the selected tenant and deployed to it when you push the template to a specific site.

Step 3 Create a DHCP Option Policy.

- a) From the **+Create Object** dropdown, select **DHCP Option Policy**.
- b) In the right properties sidebar, provide the **Name** for the policy.
- c) (Optional) Click **Add Description** and provide a description for the policy.
- d) Click **Add Option**.
- e) Provide option details.

For each DHCP option, provide the following:

- **Name** – While not technically required, we recommend using the same name for the option as listed in [RFC 2132](#).

For example, `Name Server`.

- **Id** – Provide the value if the option requires one.

For example, a list of name servers available to the client for the Name Server option.

- **Data** – Provide the value if the option requires one.

For example, a list of name servers available to the client for the Name Server option.

- Click **OK** to save.
- Repeat the previous substeps for any additional options in the same DHCP Option policy.
- Repeat this step to create any additional DHCP Option policies.

Assigning DHCP Policies

This section describes how to assign a DHCP policy to a bridge domain.



Note If you make changes to the DHCP policy after you have assigned it to a bridge domain and deployed the bridge domain to one or more sites, you will need to re-deploy the bridge domain it for the DHCP policy changes to be updated on each site's APIC.

Before you begin

You must have the following already configured:

- A DHCP relay policy, as described in [Creating DHCP Relay Policies, on page 194](#).
- (Optional) A DHCP option policy, as described in [Creating DHCP Option Policies, on page 196](#).
- The bridge domain to which you will assign the DHCP policy, as described in the [Creating Schemas and Templates, on page 29](#) chapter.

Procedure

- Step 1** Log in to your Nexus Dashboard Orchestrator GUI.
- Step 2** From the left navigation menu, select **Application Management > Schemas**.
- Step 3** Select the schema where the bridge domain is defined.
- Step 4** Scroll down to the **Bridge Domain** area and select the bridge domain.
- Step 5** In the right sidebar, scroll down and check the **DHCP Policy** option checkbox.
- Step 6** From the **DHCP Relay Policy** dropdown, select the DHCP policy you want to assign to this BD.
- Step 7** (Optional) From the **DHCP Option Policy** dropdown, select the option policy.

A DHCP option policy provides additional options to be passed to the DHCP relay. For additional details see [Creating DHCP Option Policies, on page 196](#).

Step 8 Assign the bridge domain to any EPG that needs access to the DHCP server via the relay.

Creating DHCP Relay Contract

DHCP packets are not filtered by contracts but contracts are required in many cases to propagate routing information within the VRF and across VRFs. Even though the DHCP packets are not filtered it is recommended to configure contracts between the client EPG and the EPG configured as the provider in the DHCP relay policy.

This section describes how to create a contract between the EPG that contains the DHCP server and the EPG that contains endpoints that need to use the relay. Even though you have already created and assigned the DHCP policy to the bridge domain and the bridge domain to the clients' EPG, you must create and assign the contract to enable programming of routes to allow client to server communication.

Before you begin

You must have the following already configured:

- A DHCP relay policy, as described in [Creating DHCP Relay Policies, on page 194](#).
- (Optional) A DHCP option policy, as described in [Creating DHCP Option Policies, on page 196](#).
- The bridge domain to which you have assigned the DHCP policy, as described in [Assigning DHCP Policies, on page 197](#).

Procedure

- Step 1** Log in to your Nexus Dashboard Orchestrator GUI.
- Step 2** From the left navigation menu, select **Schemas**.
- Step 3** Select the schema where you want to create the contract.
- Step 4** Create a contract.

DHCP packets are not filtered by the contract so no specific filter is required, but a valid contract should be created and assigned to ensure proper BD and routes deployment.

- a) Scroll down to the **Contracts** area and click + to create a contract.
- b) In the right sidebar, provide the **Display Name** for the contract.
- c) From the **Scope** dropdown, select the appropriate scope.

Because the DHCP server EPG and application EPG must be in the same tenant, you can select one of the following:

- `vrf`, if both EPGs are in the same VRF
- `tenant`, if the EPGs are in different VRFs

- d) You can leave the **Apply Both Directions** knob on.

- Step 5** Assign the contract to the DHCP relay EPG.
- a) Browse to the template where the EPG is located.
 - b) Select the EPG or external EPG where the DHCP server resides.
This is the same EPG you selected when creating the DHCP relay policy.
 - c) In the right sidebar, click **+Contract**.
 - d) Select the contract you created and `provider` for its type.
- Step 6** Assign the contract to the application EPG whose endpoints require DHCP relay access.
- a) Browse to the template where the application EPG is located.
 - b) Select the application EPG.
 - c) In the right sidebar, click **+Contract**.
 - d) Select the contract you created and `consumer` for its type.
-

Verifying DHCP Relay Policies in APIC

This section describes how to verify that the DHCP relay policies you have created and deployed using the Nexus Dashboard Orchestrator are correctly pushed to each site's APIC. The DHCP policies you create are pushed to the APIC when the bridge domain to which the policy is associated is deployed to a site.

Procedure

- Step 1** Log in to the site's APIC GUI.
- Step 2** From the top navigation bar, select **Tenants** > **<tenant-name>**.
Select the tenant where you deployed the DHCP policy.
- Step 3** Verify that the DHCP relay policy is configured in APIC.
In the left tree view, navigate to **<tenant-name>** > **Policies** > **Protocol** > **DHCP** > **Relay Policies**. Then confirm that the DHCP relay policy you configured has been created.
- Step 4** Verify that the DHCP option policy is configured in APIC.
If you have not configured any DHCP option policies, you can skip this step.
In the left tree view, navigate to **<tenant-name>** > **Policies** > **Protocol** > **DHCP** > **Option Policies**. Then confirm that the DHCP option policy you configured has been created.
- Step 5** Verify that the DHCP policy is correctly associated with the bridge domain.
In the left tree view, navigate to **<tenant-name>** > **Networking** > **Bridge Domains** > **<bridge-domain-name>** > **DHCP Relay Labels**. Verify that the DHCP policy is also associated with the deployed bridge domain.
-

Editing or Deleting Existing DHCP Policies

This section describes how to edit or delete a DHCP relay or option policy.

**Note**

- If you make changes to the DHCP policy after you have assigned it to a bridge domain and deployed the bridge domain to one or more sites, you will need to re-deploy it for the DHCP policy changes to update on each site's APIC.
- You cannot delete policies that are associated with one or more bridge domains, you must first unassign the policy from every bridge domain.

Procedure

-
- Step 1** Log in to your Nexus Dashboard Orchestrator GUI.
- Step 2** From the left navigation menu, select **Policies**.
- Step 3** Click the actions menu next to the DHCP policy and select **Edit** or **Delete**.
-



CHAPTER 18

EPG Preferred Group

- [EPG Preferred Groups Overview and Limitations, on page 201](#)
- [Configuring EPGs for Preferred Group, on page 202](#)

EPG Preferred Groups Overview and Limitations

By default, Multi-Site architecture allows communication between EPGs only if a contract is configured between them. If there is no contract between the EPGs, any inter-EPG communication is explicitly disabled. The Preferred Group (PG) feature allows you to specify a set of EPGs that are part of the same VRF to allow full communication between them with no need for contracts to be created.

Preferred Group vs Contracts

There are two types of policy enforcements available for EPGs in a VRF which is stretched to multiple sites with a contract preferred group configured:

- **Included EPGs** – Any EPG that is a member of a preferred group can freely communicate with all other EPGs in the group without any contracts. The communication is based on the `source-any-destination-any-permit` default rule and appropriate Multi-Site translations.
- **Excluded EPGs** – EPGs that are not members of preferred groups continue to require contracts to communicate with each other. Otherwise, the default `source-any-destination-any-deny` rule applies.

The contract preferred group feature allows for greater control and ease of configuring communication between EPGs across sites in a stretched VRF context. If two or more EPGs in the stretched VRF require open communication while others must have only limited communication, you can configure a combination of a contract preferred group and contracts with filters to control the inter-EPG communication. EPGs that are excluded from the preferred group can only communicate with other EPGs if there is a contract in place to override the `source-any-destination-any-deny` default rule.

Stretched vs Shadowed

If EPGs from multiple sites are configured to be part of the same contract preferred group, the Nexus Dashboard Orchestrator creates shadows of each site's EPGs in the other sites in order to correctly translate and program the inter-site connectivity from the EPGs. Contract preferred group policy construct is then applied in each site between a real and shadow EPG for inter-EPG communication.

For example, consider a web-service EPG1 in Site1 and an app-service EPG2 in Site2 added to the contract preferred group. Then if EPG1 wants to access EPG2, it will first be translated to a shadow EPG1 in Site2

and then be able to communicate with EPG2 using the contract preferred group. Appropriate BDs are also stretched or shadowed if the EPG under it is part of a contract preferred group.

VRF Preferred Group Setting

When you configure preferred groups directly in the APIC, you have to explicitly enable the setting on the VRF first before enabling PG membership on individual EPGs. If the PG setting on the VRF is disabled, the EPGs would not be able to communicate without contracts even if they are part of that VRF's preferred group.



Note Beginning with Release 4.0(1), PG configuration in NDO follows the same approach as it does in APIC. In other words, the PG setting on the VRF must be explicitly enabled for the EPGs that are part of that VRF to use the PG configuration.

Nexus Dashboard Orchestrator releases prior to Release 4.0(1) did not allow you to manage the PG setting on VRFs in the GUI, but instead adjusts the setting dynamically as follows:

- If you create and manage the VRF from NDO, NDO will dynamically enable or disable VRF PG value based on whether any EPGs that belong to that VRF are part of the preferred group.

In other words, when you add one or more EPGs to the preferred group, NDO automatically enables the PG setting on the VRF. When you remove the last EPG from the preferred group, NDO disables the VRF flag.

- If you want to permanently enable the PG option on a VRF, you can enable PG on the VRF directly in the APIC first, then import that VRF into NDO.

NDO will preserve the setting and not disable it automatically even if you remove every EPG from the VRF's preferred group.

- If you import the VRF from APIC without first changing the PG setting, NDO will manage the object as if it was created from NDO and overwrite the PG setting dynamically based on EPG membership.

Limitations

The following guidelines and limitations apply when using EPG Preferred Groups:

- Preferred Groups are not supported for intersite L3Out external EPGs.
- EPGs and External EPGs objects in a given VRF must not be configured as part of the Preferred Group if vzAny for that VRF is already consuming or providing a contract.

Configuring EPGs for Preferred Group

This section describes how to enable the Preferred Group (PG) configuration on the VRF and the EPGs.

Before you begin

You must have one or more EPGs added to a schema template.

Procedure

-
- Step 1** Log in to your Nexus Dashboard and open the Nexus Dashboard Orchestrator service.
- Step 2** From the left navigation pane, choose **Application Management > Schemas**.
- Step 3** Enable PG on the VRF.
- Open the Schema that contains the VRF used by the EPGs you want to include in the preferred group.
 - From the **View Overview** dropdown, select the template that contains the VRF.
 - Select the VRF.
 - In the properties sidebar on the right, check the **Preferred Group** checkbox.
- This enabled the PG configuration on that VRF. You will need to enable the PG setting on 2 or more EPGs which you want to be part of the preferred group as described in the next step.
- Click **Save** to save the template changes.
- Step 4** Configure one or more EPGs to be part of the preferred group.
- Note**
You must not have a preferred group where some EPGs are managed by Nexus Dashboard Orchestrator and some are managed locally by the APIC.
- If you have an existing preferred group in any of the APICs and are planning to import the EPGs from that preferred group into Nexus Dashboard Orchestrator, you must import all EPGs in the group.
- If the EPGs you want to include in the preferred group are in a different schema or template, navigate to that template.
 - Select an EPG.
 - In the right properties bar, check the **Include in Preferred Group** checkbox.
 - Click **Save** to save the template changes.
- Step 5** (Optional) Verify that all EPGs have been added to the preferred group.
- You can view the full list of EPGs that are configured to be part of the preferred group by selecting a VRF and checking the **Preferred Group EPGs** list in the properties sidebar on the right.
-



CHAPTER 19

Intersite L3Out

- [Intersite L3Out Overview, on page 205](#)
- [Intersite L3Out Guidelines and Limitations, on page 206](#)
- [Configuring External TEP Pool, on page 207](#)
- [Creating or Importing Intersite L3Out and VRF, on page 208](#)
- [Configuring External EPG to Use Intersite L3Out, on page 210](#)
- [Creating a Contract for Intersite L3Out, on page 212](#)
- [Use Cases, on page 215](#)

Intersite L3Out Overview

Prior to Release 2.2(1), each site managed by the Nexus Dashboard Orchestrator required its own local L3Out configured in order to route traffic out of the fabric, which often resulted in lack of communication between endpoints in one site and a service (such a firewall, server load balancer, or mainframe) connected to the L3Out of another site.

Release 2.2(1) adds a feature that enables a number of scenarios in which endpoints located in one site are able to establish connectivity with entities, such as external network, mainframe, or service nodes, reachable through a remote L3Out.

These include the following:

- L3Out across sites—endpoints in an application EPG in one site using an L3Out in another site (both part of the same VRF).
- Intersite transit routing—establishing communication between entities (such as endpoints, network devices, service nodes) connected behind L3Outs deployed in different sites (both L3Outs part of the same VRF).
- Shared services for intersite L3Out—application EPG to remote L3Out or intersite transit routing.

The following sections are divided into the generic GUI procedures you can follow to create the objects required to implement intersite L3Out use cases followed by overview and workflows specific to each supported use case scenario.



Note The term "intersite L3Out" refers to the functionality allowing communication to external resources reachable via the L3Out connection of a remote site. However, in this document, the term may also be used to indicate the specific remote L3Out object.

Intersite L3Out Guidelines and Limitations

When configuring intersite L3Out, you must consider the following:

- Intersite L3Out is supported for IPv4 and IPv6.
- With intersite L3Out, in addition to the BGP eVPN sessions that are always established between sites in Multi-Site topology, MP BGP VPNv4 (or VPNv6) sessions are created to support the intersite L3Out feature.
- If you are upgrading from a release prior to Release 2.2(1), any existing External EPG to L3Out association at the site-local level will be preserved. In addition, the Nexus Dashboard Orchestrator will now support creation of an L3Out and associating it with an External EPG at the template level.

When creating a new L3Out in a schema template and associating it to an existing External EPG:

- If the L3Out has the **same name** as the L3Out already defined in the APIC, the Orchestrator will take ownership of that L3Out but will not manage the configuration of L3Out node profiles, interface profiles, protocol settings, or route control settings.



Note If the L3Out already exists in APIC, we recommend importing it into Nexus Dashboard Orchestrator along with any associated external EPG instead of creating a new L3Out with the same name from NDO.

If you then choose to delete this L3Out from the Orchestrator, it will no longer be managed by the Orchestrator, but any previously existing L3Out configuration will be preserved in the APIC.

- If the L3Out has a **different name** than the APIC defined L3Out the external EPG will be removed from the APIC defined L3Out and added to the L3Out defined in the Orchestrator. If this is the only external EPG under the APIC defined L3Out this can cause the configuration to be removed from the border leaves and can impact traffic.
 - If you choose to downgrade to a release prior to Release 2.2(1), the L3Outs created in the Orchestrator NDO will no longer exist in the template so any template-level association between External EPG and L3Out will be removed. In this case, you will need to manually re-configure the External EPG to L3Out association at the site-local level. Any site-local associations will be preserved during the downgrade.
 - You can now associate a bridge domain in one site with the L3Out in another site, however they must both be in the same VRF.
- This association is performed at the site-local level and is required to advertise the BD subnet out of the remote L3Out and ensure that inbound traffic to the BD can be maintained even if the local L3Out failed.
- The Policy Control Enforcement direction for the VRF associated to the intersite L3Out must be kept configured in the default ingress mode.

- The following scenarios are not supported with intersite L3Out and remote leaf (RL):
 - Transit routing between L3Outs deployed on RL pairs associated to separate sites
 - Endpoints connected to a RL pair associated to a site communicating with the L3Out deployed on the RL pair associated to a remote site
 - Endpoints connected to the local site communicating with the L3Out deployed on the RL pair associated to a remote site
 - Endpoints connected to a RL pair associated to a site communicating with the L3Out deployed on a remote site
- The following other features are not supported with intersite L3Out in Multi-Site:
 - Multicast receivers in a site receiving multicast from an external source via another site L3Out. Multicast received in a site from an external source is never sent to other sites. When a receiver in a site receives multicast from an external source it must be received on a local L3Out.
 - An internal multicast source sending multicast to an external receiver with PIM-SM any source multicast (ASM). An internal multicast source must be able to reach an external Rendezvous Point (RP) from a local L3Out
 - GOLF
 - Preferred Groups for External EPG

Configuring External TEP Pool

Intersite L3Out requires an external TEP address for the border leaf switches in each pod. If you already have an external TEP pool configured, for example for another feature such as Remote Leaf, the same pool can be used. The existing TEP pool will be inherited by the Nexus Dashboard Orchestrator and shown in the GUI as part of the infra configuration. Otherwise, you can add a TEP pool in the GUI, as described in this section.



Note Every pod must be assigned a unique TEP pool and it must not overlap with any other TEP pool in the fabric

Procedure

- Step 1** Log in to your Nexus Dashboard Orchestrator.
- Step 2** In the left navigation menu, select **Infrastructure > Site Connectivity**.
- Step 3** In the top right of the main pane, click **Configure**.
- Step 4** In the left sidebar, select the site you want to configure.
- Step 5** In the main window, click a pod in the site.
- Step 6** In the right sidebar, click **+Add TEP Pool**.
- Step 7** In the **Add TEP Pool** window, specify the external TEP pool you want to configure for that site.

Note

You must ensure that the TEP pool you are adding does not overlap with any other TEP pools or fabric addresses.

Step 8 Repeat the process for each site and pod where you plan to use intersite L3Outs.

Creating or Importing Intersite L3Out and VRF

This section describes how to create an L3Out and associate it to a VRF in the Orchestrator GUI, which will then be pushed out to the APIC site, or import an existing L3Out from one of your APIC sites. You will then associate this L3Out with an external EPG and use that external EPG to configure specific intersite L3Out use cases.



Note The VRF you assign to the L3Out can be in any template or schema, but it must be in the same tenant as the L3Out.

Procedure

Step 1 Log in to your Nexus Dashboard Orchestrator.

Step 2 From the left navigation pane, select **Application Management > Schemas**.

Step 3 Select the schema and then the template where you want to create or import the VRF and L3Out.

We recommend creating the L3Out in a template that is associated with a single site, in which case the L3Out will be created in that site only.

Alternatively, you can choose to create the L3Out in a template that is associated to multiple sites. In this case the L3Out will be created with the same name across all sites, which may bring some functional restrictions, as explained later in this chapter

Step 4 Create a new VRF and L3Out.

If you want to import an existing L3Out, skip this step.

Note

While you can create the L3Out object in the Orchestrator and push it out to the APIC, the physical configuration of the L3Out must be done in the APIC.

a) Scroll down to the **VRF** area and click the + icon to add a new VRF.

If you already have the VRF you plan to use for the L3Out, skip this substep.

In the right sidebar, provide the name for the VRF, for example `vrf-l3out`

b) Scroll down to the **L3Out** area and click the + icon to add a new L3Out.

In the right sidebar, provide the required information.

c) Provide the name for the L3Out, for example `l3out-intersite`.

d) From the **Virtual Routing & Forwarding** dropdown, select the VRF.

Select the VRF you created in the first substep or choose a previously existing VRF.

Step 5 Import an existing VRF and L3Out.

If you created a new L3Out in previous step, skip this step.

Click **Import** in the main window pane to open at the

POLICY TYPE	SELECT TO IMPORT	INCLUDE RELATIONS
APPLICATION PROFILE 0 out of 1	☒ 1 VRF	☐
EPG 0 out of 7		
EXTERNAL EPG 0 out of 1		
CONTRACT 0 out of 5		
FILTER 0 out of 2		
VRF 1 out of 3		
BD 0 out of 15		
L3OUT 1 out of 1		

Import

- At the top of the main template view, click **Import**.
- Select the site from which you want to import the L3Out.
- In the import window's **Policy Type** menu, select **L3Out**.
- Check the L3Out you want to import.

By default, importing the L3Out will also import the corresponding VRF. This may not be desirable when importing the L3Out in a site specific template as you would typically define the VRF in a stretched template associated to multiple sites. In this case, disable the **Include Relations** option before importing the L3Out. In this case, you will also need to re-map the L3Out to the correct VRF after importing it.

- Click **Import**.

- f) If you imported only the L3Out, select it in the template view and associate it to the appropriate VRF.
-

Configuring External EPG to Use Intersite L3Out

This section describes how to create an external EPG that will be associated to the intersite L3Out. You can then use this external EPG and contracts to configure specific use cases for endpoints in one site to use an L3Out in another site.

Before you begin

Create the L3Out and associate it with a VRF as described in [Creating or Importing Intersite L3Out and VRF, on page 208](#).

Procedure

Step 1 Select the template where you want to create the external EPG.

If you create the external EPG in a template that is associated to multiple sites, the external EPG will be created on all of those sites. This is recommended when the external EPG's L3Outs provide access to a set of common external resources, for example the WAN.

If you create the external EPG in a template that is associated with a single site, the external EPG will be created in that site only. This is recommended when the external EPG's L3Out provides access to external resources accessible only from that site.

Step 2 Scroll down to the **External EPG** area and click the + icon to add an external EPG.

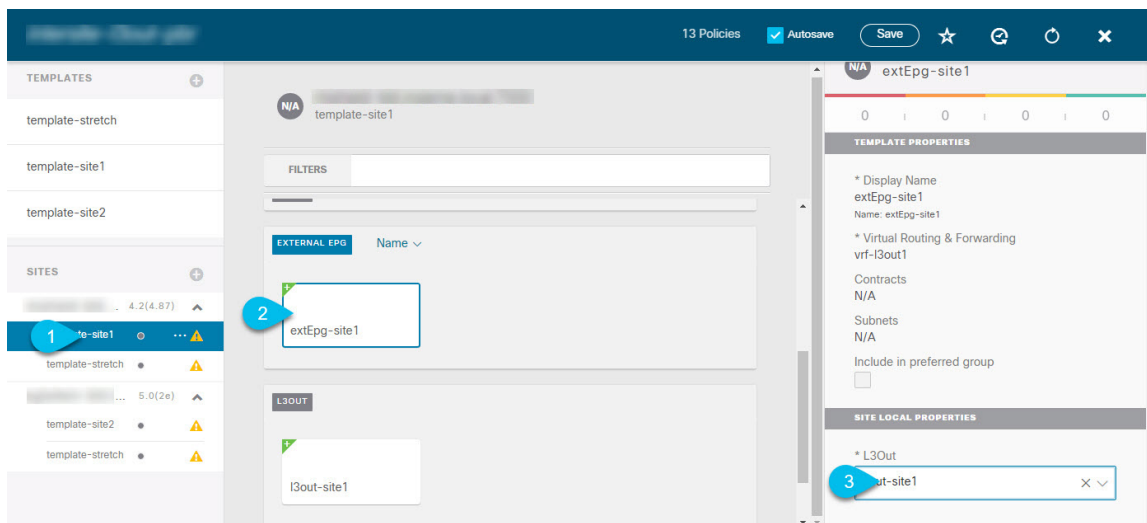
In the right sidebar, provide the required information.

- a) Provide the name for the external EPG, for example `eepg-intersite-l3out`.
- b) From the **Virtual Routing & Forwarding** dropdown, select the VRF you created and used for the L3Out.

Step 3 Map the external EPG to the L3Out.

You can map the external EPG to an L3Out at the site level or at the template level. We recommend creating the mapping at the site level because commonly each site defines a local L3Out with a unique name so the external EPG can be selectively mapped to each site specific L3Out independent of whether the external EPG itself is stretched.

To associate an L3Out with the external EPG at the site-local level:

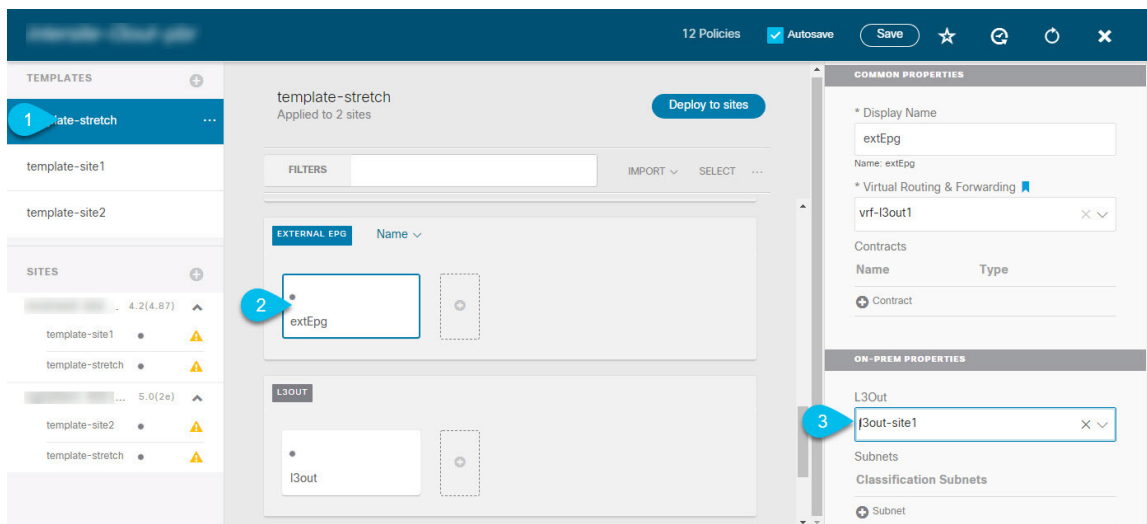


- In the left sidebar of the schema view, select the site where the external EPG is deployed.
- Scroll down to the **External EPG** area and select the external EPG.
- In the right sidebar, scroll down to the **L3Out** dropdown and choose the intersite L3Out you created.

In this case, both the APIC-managed and the Orchestrator-managed L3Outs will be available for selection. You can select either the L3Out you have created in the previous section specifically for this or pick an L3Out that exists in the site's APIC.

Alternatively, you can map the external EPG to an L3Out at the template level. While this could ease the configuration in deployments where multiple sites have defined the same L3Out name, we do not recommend this approach as it allows less flexibility for the type of connectivity that can be established between the fabrics that are part of the Multi-Site domain and the external routed network. For example, it would not be possible to control where a specific BD's subnets are advertised because mapping the BD to the L3Out would cause the BD subnet to be advertised out of all the L3Outs in all the sites since all the L3Outs have the same name.

To associate an L3Out with the external EPG at the template level:



- In the left sidebar of the schema view, select the template where the external EPG is located
- Scroll down to the **External EPG** area and select the external EPG.

- c) In the right sidebar, scroll down to the **L3Out** dropdown and choose the intersite L3Out you created.

In addition, it is possible to migrate the configuration of an external EPG initially associated to the L3Out at the template level to a site-level mapping by removing the VRF association on the external EPG, re-associating the external EPG to the same VRF, then mapping the L3Outs at the site level. If this process is completed at once before deploying the template, there would be no traffic impact when pushing the new configuration as no changes are actually applied on the APIC side.

Step 4 Configure one or more subnets for the external EPG.

- Select the external EPG.
- In the right sidebar, click **+Add Subnet**.
- In the **Add Subnet** window, provide the classification subnet and the required options.

The prefixes and options you configure depend on the specific use cases:

- To classify the inbound traffic as belonging to the external EPG, select the **External Subnets for External EPG** flag for the specified prefix. Depending on the specific use case, this allows you to apply a contract with an internal EPG or with the external network domain reachable via a remote L3Out.
- To advertise the external prefixes learned from another L3Out (in the same site or in a remote site) out of this L3Out, select the **Export Route Control** flag for the specified prefix. When specifying the `0.0.0.0/0` prefix, the **Aggregate Export** flag can be selected to advertise all prefixes out of the L3Out; if the **Aggregate Export** flag is not enabled, only the default route `0.0.0.0/0` would be advertised, if present in the routing table of the border leaf nodes.
- To filter out specific routes received from the external network, select the **Import Route Control** flag for the specified prefix. If specifying the `0.0.0.0/0`, you can also choose the **Aggregate Import** option.

Note that this is possible only when peering BGP with the external routers.

- To leak routes to different VRFs, select the **Shared Route Control** and the associated **Aggregate Shared Routes** flags, as well as the **Shared Security Import** flag. These options are required for the specific use case of inter-VRF shared L3Out and inter-VRF intersite transit routing.

Creating a Contract for Intersite L3Out

This section describes how to create a filter and a contract you will use to enable communication between an application EPG deployed in a site and the external EPG associated to an L3Out in a different site (intersite L3Out functionality).

Procedure

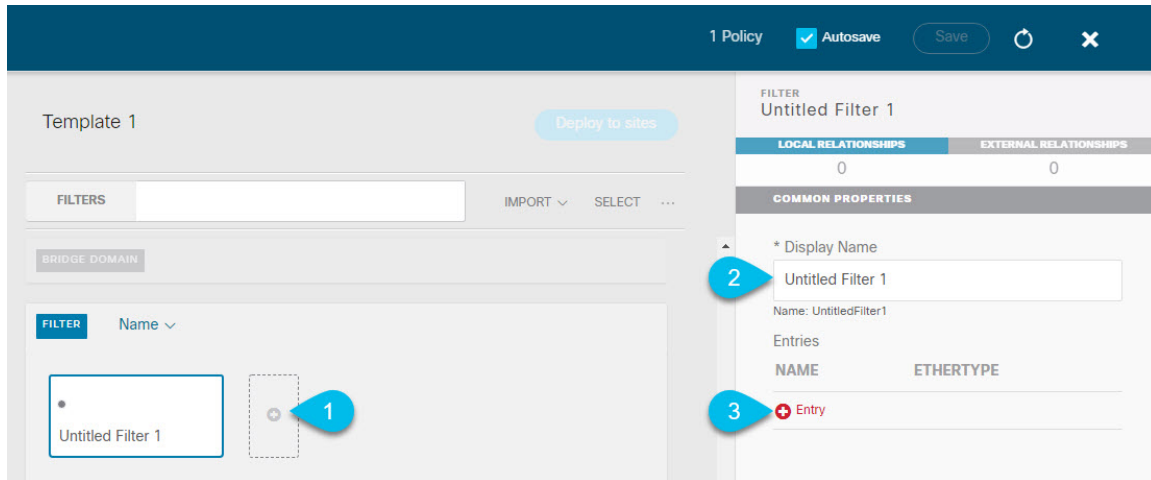
Step 1 Select the template where you want to create contract and filter.

You can use the same schema and template where you created the L3Out, VRF, and the external EPG or you can choose a different schema and template.

Because the contract is applied to objects (EPGs and external EPGs) deployed in different sites, we recommend defining it in a template associated to multiple sites. However, this is not strictly required and even if the contract and filters are

defined only as local objects in Site1, NDO will create the corresponding shadow objects in a remote Site2 when a local EPG or external EPG in Site2 needs to consume or provide that contract.

Step 2 Create a filter.



- In the middle pane, scroll down to the **Filter** area, then click + to create a filter.
- In the right pane, provide the **Display Name** for the filter.
- In the right pane, click + **Entry**.

Step 3 Provide the filter details.

Add Entry

COMMON PROPERTIES

Name

icmp

Description

Ether Type

ip

IP Protocol

icmp

Destination port range from

unspecified

Destination port range to

unspecified

ON-PREM PROPERTIES

☐ Match only fragments

☐ stateful

ARP flag

unspecified

Source port range from

unspecified

Source port range to

unspecified

TCP session rules

unspecified

4

Save

- a) Provide the **Name** for the filter.
- b) Choose the **Ether Type**.

214

Cisco Nexus Dashboard Orchestrator Configuration Guide for ACI Fabrics, Release 4.0(x)

For example, `ip`.

- c) Choose the **IP Protocol**.

For example, `icmp`.

- d) Leave other properties unspecified.
- e) Click **Save** to save the filter.

Step 4 Create a contract

- a) In the middle pane, scroll down to the **Contract** area and click **+** to create a contract.
- b) In the right pane, provide the **Display Name** for the contract
- c) Select the appropriate **Scope** for the contract.

If you plan to configure different VRFs for the intersite L3Out and application EPG, you must select `tenant` for the scope. Otherwise, if both are in the same VRF, you can set the scope to `vrf`.

- d) Toggle the **Apply both directions** knob if you want the same filter to apply for both consumer-to-provider and provider-to-consumer directions.

If you enable this option, you will need to provide the filters only once and they will apply for traffic in both directions. If you leave this option disabled, you will need to provide two sets of filter chains, one for each direction.

Step 5 Assign the filters to the contract

- a) In the right pane, scroll down to the **Filter Chain** area and click **+** **Filter** to add a filter to the contract.

If you disabled the `Apply both directions` option, repeat this step for the other filter chain.

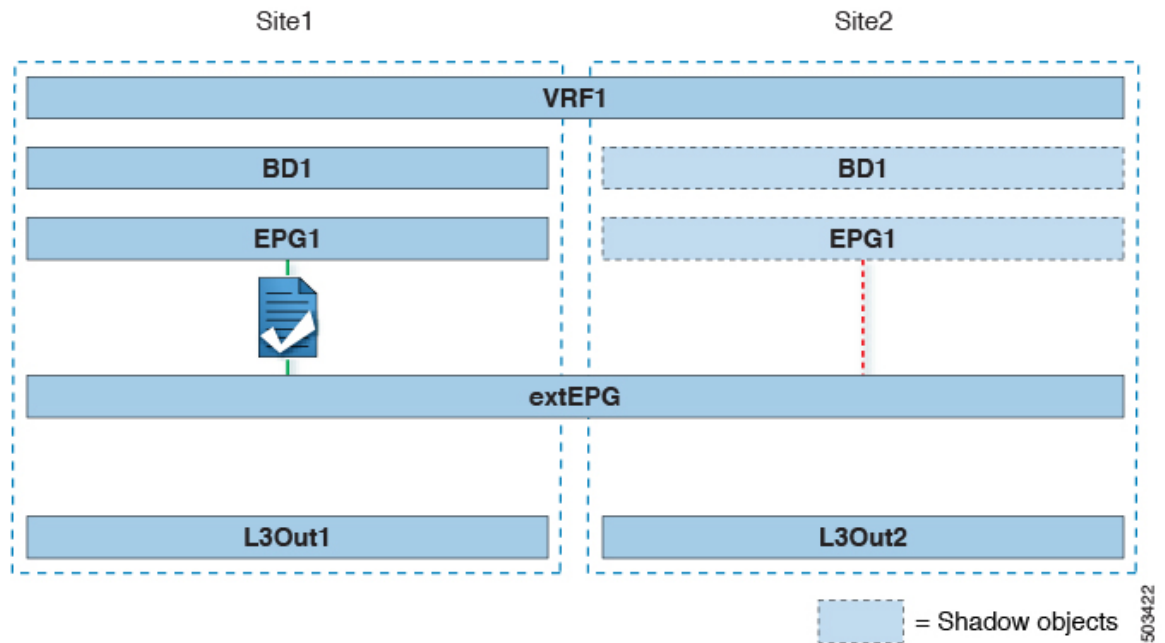
- b) In the **Add Filter Chain** window that opens, select the filter you added in previous step from the **Name** dropdown menu.
- c) Click **Save** to add the filter to the contract.

Use Cases

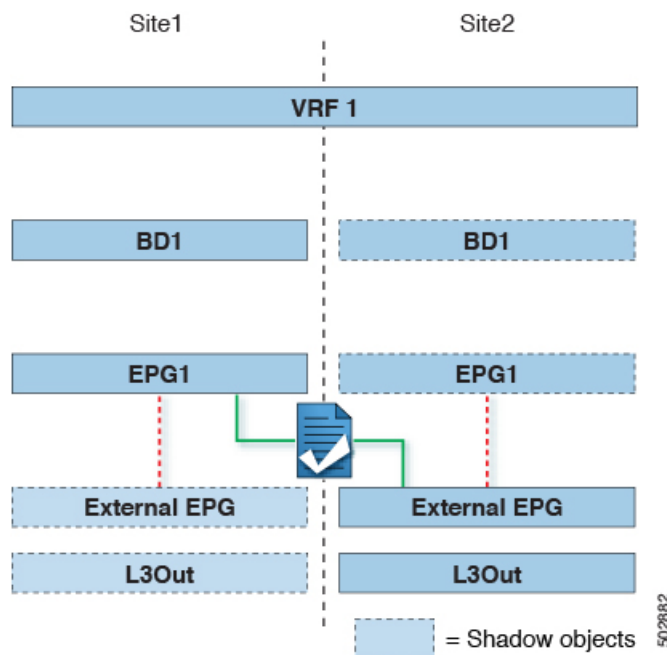
Intersite L3Out for Application EPGs (Intra-VRF)

This section describes the configuration required to allow endpoints that are part of an application EPG to communicate with the external network domain reachable through an L3Out deployed in another site but within the same VRF (intra-VRF).

The first figure below shows a stretched external EPG and the associated L3Outs which will be created in both sites. An application EPG (`EPG1`) is created in Site 1 and has a contract with the external EPG. This use case is recommended when the L3Outs in the separate sites provide access to a common set of external resources. It simplifies the policy definition and external traffic classification, while still allowing you to apply route-map policies separately on each L3Out for the independent APIC domains.

Figure 14: Stretched External EPG

The second figure below shows a similar use case but with the external EPG being deployed to only the site where the physical L3Out is located. The application EPG and the contract are configured in the same exact way to allow the traffic flow between the EPG in one site and the physical L3Out in the other.

Figure 15: Non-Stretched (Site-Local) External EPG

The following steps describe the configuration required to implement the use case shown in Figure 1, which represents the most common scenario. If you want to deploy the use case shown in Figure 2, you can adapt the procedure with minor changes.

Before you begin

You need to have the following already configured:

- A schema with three templates.

Create a template for each site (for example, `template-site1` and `template-site2`) where you will configure the objects unique to that site, such as the application EPG and the L3Outs. In addition, create a separate templates (for example, `template-stretched`) that you will use for the stretched objects, which in this case will be the external EPG.

- The L3Outs in each site, as described in the [Creating or Importing Intersite L3Out and VRF, on page 208](#) section.

In this use case, a separate L3Out will be imported or created in each site-specific template.

- The external EPG for the intersite L3Out, as described in [Configuring External EPG to Use Intersite L3Out, on page 210](#).

In this use case, the external EPG is configured as a stretched object that is defined in the stretched template (`template-stretched`). Assuming that the external EPG provides access to the entire external address space, we recommend configuring a `0.0.0.0/0` prefix for classification to avoid specifying a long list of more specific prefixes.

- The contract you will use between the application EPG and the L3Out external EPG, as described in [Creating a Contract for Intersite L3Out, on page 212](#).

We recommend creating the contract and the filter in the stretched template (`template-stretched`).

Procedure

-
- | | |
|---------------|--|
| Step 1 | Log in to your Nexus Dashboard Orchestrator. |
| Step 2 | From the left navigation pane, select Application Management > Schemas . |
| Step 3 | Select the schema and template for the application EPG and bridge domain.

In this use case, you will associate the template to Site1. |
| Step 4 | Configure an application EPG and its bridge domain belonging to the same VRF as the L3Out.

If you already have an EPG that will use the intersite L3Out, you can skip this step.

You can create a new or import an existing EPG and bridge domain as you typically would. |
| Step 5 | Assign the contract to the application EPG.
a) Select the EPG.
b) In the right sidebar, click +Contract .
c) Select the contract you created in previous section and its type.

You can choose whether the application EPG is the <code>consumer</code> or the <code>provider</code> . |
| Step 6 | Assign the contract to the external EPG mapped to the remote L3Out.
a) Select the <code>template-stretched</code> where the external EPG is located.
b) Select the external EPG.
c) In the right sidebar, click +Contract . |

- d) Select the contract you created in previous section and its type.

If you chose the application EPG to be the `consumer`, choose `provider` for the external EPG. Otherwise, choose `consumer` for the external EPG.

Step 7 Associate the application EPG's bridge domain with the L3Out.

This enables the BD subnet to be advertised out of the L3Out toward the external network domain. Note that the subnet(s) associated to the BD must be configured with the **Advertised Externally** option to be advertised out of the L3Out

- a) In the left sidebar, under **Sites**, select the application EPG's template.
- b) Select the bridge domain associated with the application EPG.
- c) In the right sidebar, click **+L3Out**.
- d) Select the intersite L3Out you created.

For the use case shown in Figure 1, associate the BD to both the L3Outs defined in Site1 and Site2 to ensure that the external network can have access to the EPG from both paths. Specific policies can be associated to the L3Out or to the external routers to ensure that a specific L3Out path is normally preferred for inbound traffic. We recommend this when the EPG and BD are local to a site (as in the specific example) to avoid suboptimal inbound traffic path via the remote site's L3Out.

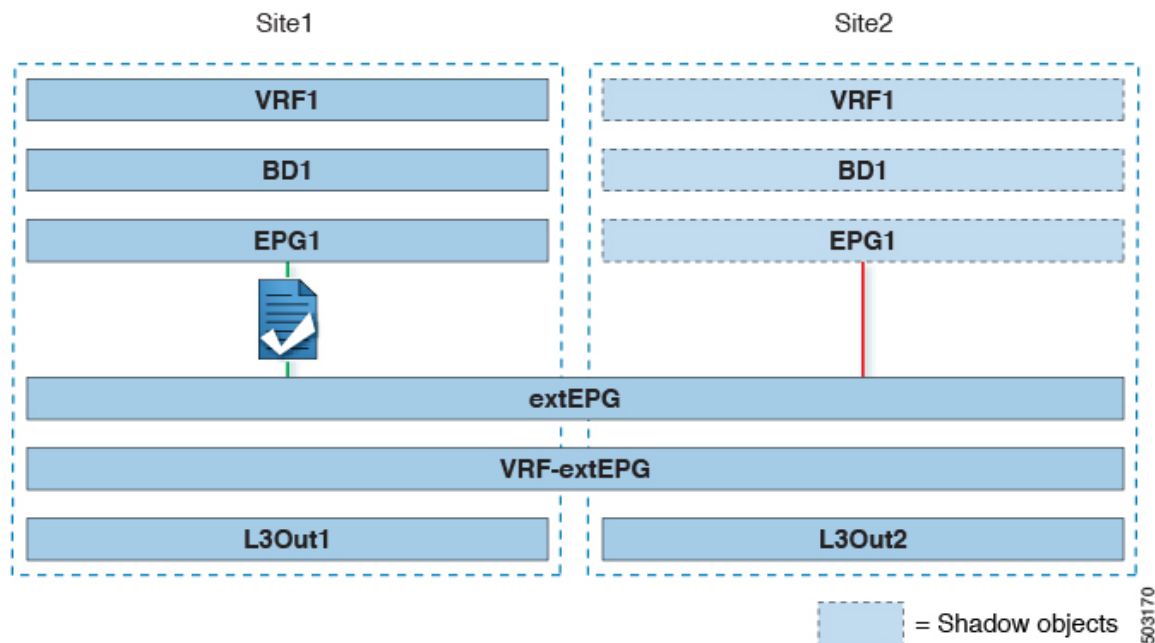
Step 8 Deploy the schema.

Shared Services with Intersite L3Out for Application EPGs (Inter-VRF)

This section describes the configuration required to allow endpoints that are part of an application EPG in one VRF to communicate with the external network domain reachable through an L3Out deployed in another site and different VRF, this is also known as "Shared Services".

This scenario is recommended when the L3Outs in separate sites provide access to a common set of external resources. It simplifies the policy definition and external traffic classification, while still allowing you to apply route-map policies separately on each L3Out for the independent APIC domains.

Figure 16: Stretched External EPG, Site-Local L3Outs and Application EPGs



The following steps describe the configuration required to implement the use case shown in Figure 3.

Before you begin

You need to have the following already configured:

- A schema with three templates.

Create a template for each site (for example, `template-site1` and `template-site2`) where you will configure the objects unique to that site, such as the application EPGs and the L3Outs. In addition, create a separate templates (for example, `template-stretched`) that you will use for the stretched objects, which in this case will be the external EPG.

- The L3Outs in each site, as described in the [Creating or Importing Intersite L3Out and VRF](#), on page 208 section.

In this use case, a separate L3Out will be imported or created in each site-specific template.

- The external EPG for the intersite L3Out, as described in [Configuring External EPG to Use Intersite L3Out](#), on page 210.

In this use case, the external EPG is configured as a stretched object that is defined in the stretched template (`template-stretched`). Assuming that the external EPG provides access to the entire external address space, we recommend configuring a `0.0.0.0/0` prefix for classification to avoid specifying a long list of more specific prefixes.

For this specific shared services use case, you are also required to enable the **Shared Route Control** and the **Shared Security Import** flags for the subnet(s) associated to the external EPG(s) of the remote L3Out. If you are using the `0.0.0.0/0` prefix for classification on the external EPG, in addition to the **Shared Route Control** flag, also enable the **Aggregate Shared Routes** flag.

- The contract you will use between the application EPG and the L3Out external EPG, as described in [Creating a Contract for Intersite L3Out](#), on page 212.

We recommend creating the contract and the filter in the stretched template (`template-stretched`).

Procedure

Step 1 Log in to your Nexus Dashboard Orchestrator.

Step 2 From the left navigation pane, select **Application Management > Schemas**.

Step 3 Select the schema and template for the application EPG and bridge domain.

In this use case, you will associate the template to Site1.

Step 4 Configure an application EPG and its bridge domain belonging to a separate VRF from the L3Out's.

If you already have an EPG that will use the intersite L3Out, you can skip this step.

You can create a new or import an existing EPG and bridge domain as you typically would.

Step 5 Assign the contract to the application EPG.

- a) Select the EPG.
- b) In the right sidebar, click **+Contract**.
- c) Select the contract you created in previous section and its type.

You can choose whether the application EPG is the `consumer` or the `provider`.

Note

If the application EPG is configured as `provider`, you need to configure the subnet already defined under the BD also under the EPG in order to leak that route into the L3Out VRF. The same flags used under the BD for the subnet should also be set under the EPG. In addition to that, for the subnet under the EPG the flag **No default SVI Gateway** should also be enabled, since the default gateway function is enabled at the BD level.

Step 6 Assign the contract to the external EPG mapped to the L3Outs.

- a) Select the `template-stretched` where the external EPG is located.
- b) Select the external EPG.
- c) In the right sidebar, click **+Contract**.
- d) Select the contract you created in previous section and its type.

If you chose the application EPG to be the `consumer`, choose `provider` for the external EPG. Otherwise, choose `consumer` for the external EPG.

Step 7 Associate the application EPG's bridge domain with the L3Out.

This enables the BD subnet to be advertised out of the L3Out toward the external network domain. Note that the subnet(s) associated to the BD must be configured with the **Advertised Externally** option to be advertised out of the L3Out

- a) In the left sidebar, under **Sites**, select the application EPG's template.
- b) Select the bridge domain associated with the application EPG.
- c) In the right sidebar, click **+L3Out**.
- d) Select the intersite L3Out you created.

For the use case shown in Figure 1, associate the BD to both the L3Outs defined in Site1 and Site2 to ensure that the external network can have access to the EPG from both paths. Specific policies can be associated to the L3Out or to the external routers to ensure that a specific L3Out path is normally preferred for inbound traffic. We recommend

this when the EPG and BD are local to a site (as in the specific example) to avoid suboptimal inbound traffic path via the remote site's L3Out.

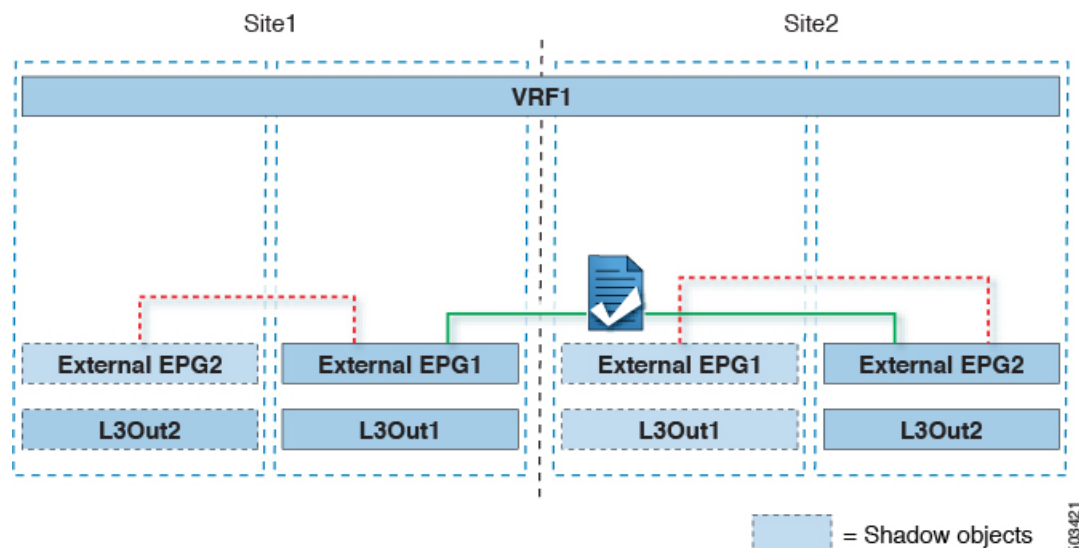
Step 8 Deploy the schema.

Intersite Transit Routing

This section describe the use cases where the Multi-Site domain acts as a distributed router allowing communication between entities (endpoints, network devices, service nodes, etc.) connected behind L3Outs deployed in different sites, a functionality normally referred to as intersite transit routing. The intersite transit routing is supported for intra-VRF as well as inter-VRF use cases.

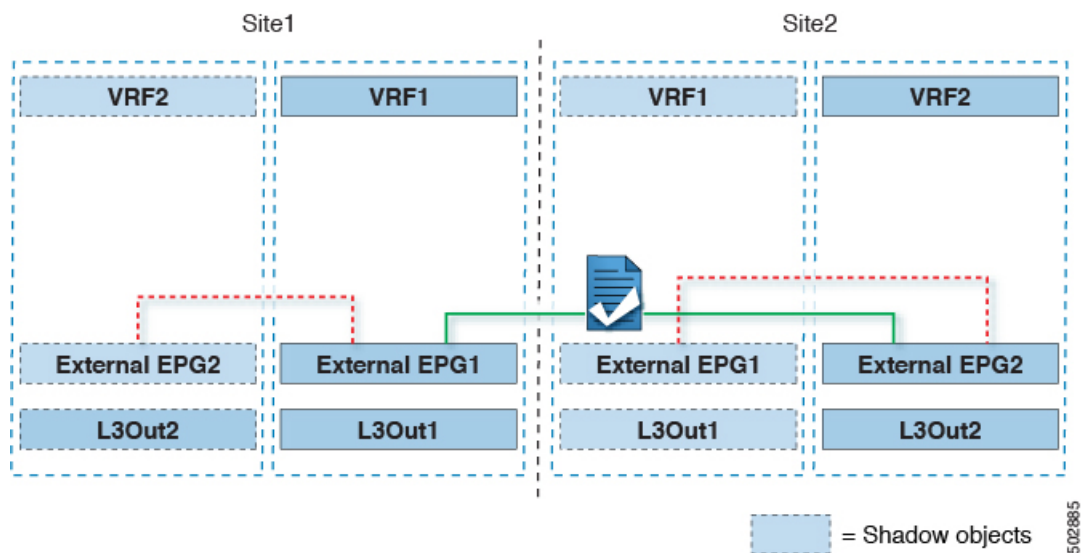
The figure below shows two L3Outs (`L3Out1` and `L3Out2`) configured in different sites. Each L3Out is associated with a respective external EPG (`External EPG1` and `External EPG2`). A contract between the two external EPGs allows communication between entities connected behind two different L3Outs in two different sites.

Figure 17: Intra-VRF Intersite Transit Routing



A similar configuration can be used when each site's L3Outs are in different VRFs.

Figure 18: Inter-VRF Intersite Transit Routing



The figures above show the two scenarios where the external EPGs and associated L3Outs are deployed as site-local objects; intersite transit routing can support all the combinations where neither external EPG is stretched, one of them is stretched, or both are stretched between sites.

When deploying intersite transit routing, the assumption is that the different external EPGs defined across sites are providing access to different external address spaces (obviously not overlapping). A couple of options are hence possible for the configuration of the prefix used for classification:

- Define the same `0.0.0.0/0` prefix on both external EPGs to ensure that inbound traffic received on the border leaf nodes of L3Out1 gets mapped to Ext-EPG1, whereas inbound traffic received on L3Out2 gets mapped to Ext-EPG2. Because the L3Outs are defined in separate fabrics, there are no conflict issues with this configuration.

The external prefixes received on L3Out1 must be advertised out of L3Out2 and vice versa. If you are using `0.0.0.0/0` as classification subnet on both external EPGs, it is sufficient to enable the **Export Route Control** and the **Aggregate Export** flags.

- Define specific prefixes for each external EPG. In this case, you must ensure that the prefixes are not overlapping to avoid a fault from being raised by the site's APIC when the shadow external EPG is created in that site for a contract between the local and remote external EPGs.

When using specific prefixes, the same prefixes configured for classification on External EPG1 must be configured with the **Export Route Control** flag set on External EPG2 and vice versa.



Note No matter which of the two classification approaches you deploy, for the inter-VRF scenario you must also set the **Shared Route Control** (in addition to **Aggregate Shared Routes** if using `0.0.0.0/0`) and the **Shared Security Import** flags.

Before you begin

You need to have the following already configured:

- A schema with three templates.

Create a template for each site (for example, `template-site1` and `template-site2`) where you will configure the objects unique to that site, such as the application EPGs and the L3Outs. In addition, create a separate templates (for example, `template-stretched`) that you will use for the stretched objects, which in this case will be the external EPG.

- The L3Outs in each site, as described in the [Creating or Importing Intersite L3Out and VRF, on page 208](#) section.

In this use case, a separate L3Out will be imported or created in each site-specific template.

- Two different external EPGs for two different L3Outs in different sites. You can use the same procedure to create both external EPGs, as described in [Configuring External EPG to Use Intersite L3Out, on page 210](#).
- The contract you will use between the L3Out external EPGs defined in each site, as described in [Creating a Contract for Intersite L3Out, on page 212](#).

We recommend creating the contract and the filter in the stretched template (`template-stretched`).

Procedure

-
- Step 1** Log in to your Nexus Dashboard Orchestrator.
- Step 2** From the left navigation pane, select **Application Management > Schemas**.
- Step 3** Assign the contract to one of the external EPGs.
- Select the schema and template where the external EPG is located.
 - Select the external EPG.
 - In the right sidebar, click **+Contract**.
 - Select the contract you created in previous section and its type.
- Choose `consumer` or `provider`.
- Step 4** Assign the contract to the other external EPG.
- Select the schema and template where the external EPG is located.
 - Browse to the template where the external EPG is located.
 - Select the external EPG.
 - In the right sidebar, click **+Contract**.
 - Select the contract you created in previous section and its type.
- Choose `provider` or `consumer`.
- Step 5** Deploy the templates to appropriate sites.
-



CHAPTER 20

Intersite L3Out with PBR

- [Intersite L3Out with PBR, on page 225](#)
- [Supported Use Cases, on page 226](#)
- [Guidelines and Limitations, on page 230](#)
- [Configuring APIC Sites, on page 230](#)
- [Creating Templates, on page 234](#)
- [Configuring Service Graph, on page 236](#)
- [Creating Filter and Contract, on page 238](#)
- [Creating Application EPG, on page 244](#)
- [Creating L3Out External EPG, on page 247](#)

Intersite L3Out with PBR

Cisco Application Centric Infrastructure (ACI) policy-based redirect (PBR) enables traffic redirection for service appliances, such as firewalls or load balancers, and intrusion prevention system (IPS). Typical use cases include provisioning service appliances that can be pooled, tailored to application profiles, scaled easily, and have reduced exposure to service outages. PBR simplifies the insertion of service appliances by using contract between the consumer and provider endpoint groups even if they are all in the same virtual routing and forwarding (VRF) instance.

PBR deployment consists of configuring a route redirect policy and a cluster redirect policy, and creating a service graph template that uses these policies. After the service graph template is deployed, you can attach it to a contract between EPGs so that all traffic following that contract is redirected to the service graph devices based on the PBR policies you have created. Effectively, this allows you to choose which type of traffic between the same two EPGs is redirected to the L4-L7 device, and which is allowed directly.

More in-depth information specific to services graphs and PBR is available in the [Cisco APIC Layer 4 to Layer 7 Services Deployment Guide](#)

PBR Support in Multi-Site Deployments

Cisco Multi-Site has supported EPG-to-EPG (east-west) and L3Out-to-EPG (north-south) contracts with PBR since Cisco APIC, Release 3.2(1). However, the L3Out-to-EPG across sites (traffic from an external endpoint in `site1` to an endpoint in `site2`) case was supported only if both sites had local L3Outs. The intersite L3Out use cases were limited to the examples and configurations described in the [Intersite L3Out, on page 205](#) chapter. Similarly, the Service Graph integration with PBR but no intersite L3Out is described in great detail in the [Cisco Multi-Site and Service Node Integration White Paper](#).

Starting with Cisco APIC, Release 4.2(5), the L3Out-to-EPG with PBR across sites (intersite L3Out) use case has been extended to support cases where the application EPG has no local L3Out or the local L3Out is down.

Supported Use Cases

The following diagrams illustrate the traffic flows between an ACI internal endpoint in application EPG and an external endpoint through the L3Out in another site in the supported intersite L3Out with PBR use cases.

The workflow to configure these examples is the same, with the only differences being whether you create the objects in the same or different VRFs (inter-VRF vs intra-VRF) and where you deploy the objects (stretched vs non-stretched):

1. Create the L4-L7 devices directly in the site's APIC, as described in [Creating and Configuring L4-L7 Devices and PBR Policies](#), on page 231.

You cannot create the devices and PBR policies from the Nexus Dashboard Orchestrator, so you will need to log in to each site's APIC directly to configure those options.

2. Create the required templates, as described in [Creating Templates](#), on page 234.

We recommend creating a single stretched template that will contain all the objects deployed to all sites. Then an extra template for each site with the objects specific to that site only.

3. Create and configure the service graph, as described in [Configuring Service Graph](#), on page 236.
4. Create the contract and filter you will use for all traffic between the application EPG and the external EPG containing the L3Out in another site, as described in [Creating Filter and Contract](#), on page 238.
5. Create the application EPG with its VRF and bridge domain, as described in [Creating Application Profile and EPG](#), on page 245.

Depending on whether you plan to stretch the application EPG or not, you will create these objects in different templates. Similarly, you can choose to use the same or different VRFs for the application EPG and the L3Out.

6. Create the L3Out, as described in [Creating or Importing Intersite L3Out and VRF](#), on page 247.
7. Create the external EPG for the L3Out, as described in [Configuring External EPG to Use Intersite L3Out](#), on page 210.

Inter-VRF vs Intra-VRF

When creating and configuring the application EPG and the external EPG, you will need to provide a VRF for the application EPG's bridge domain and for the L3Out. You can choose to use the same VRF (intra-VRF) or different VRFs (inter-VRF).

When establishing a contract between the EPGs, you will need to designate one EPG as the provider and the other one as the consumer:

- When both EPGs are in the same VRF, either one can be the consumer or the provider.
- If the EPGs are in different VRFs, the external EPG must be the provider and the application EPG must be the consumer.

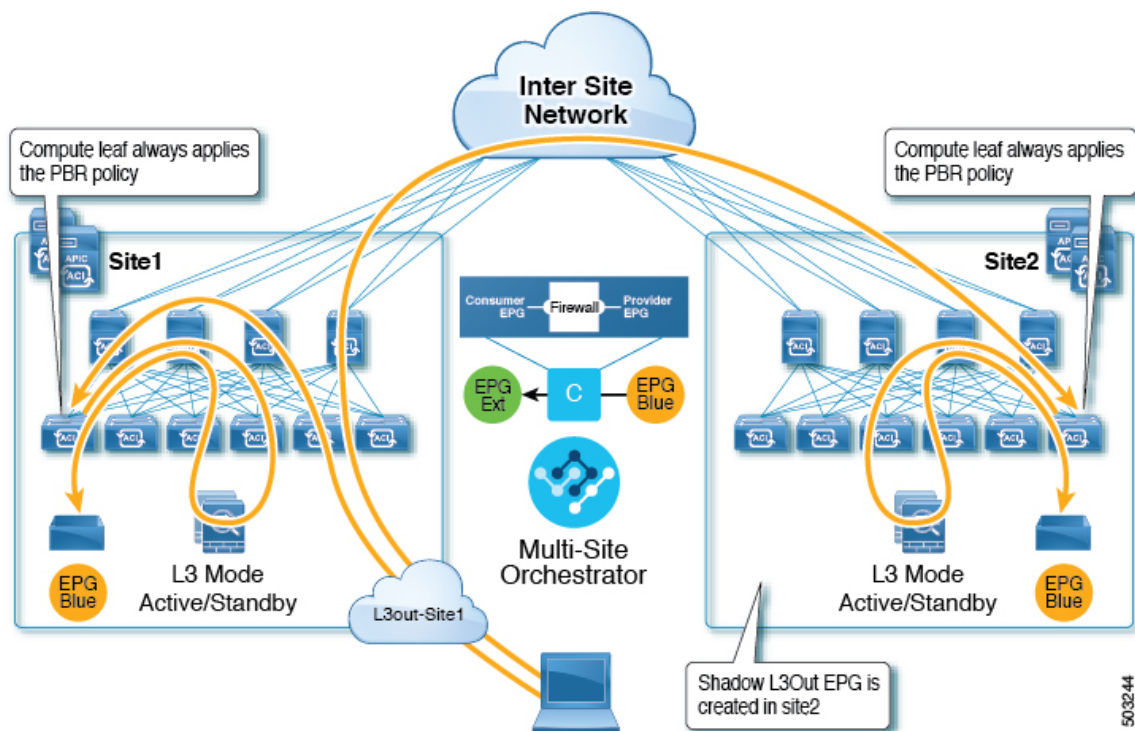
Stretched EPG

This use case illustrates a single application EPG that is stretched between two sites and a single L3Out created in only one of the sites. Regardless of whether the application EPG's endpoint is in the same site as the L3Out or the other site, traffic will go through the same L3Out. However, the traffic will always go through the service node that is local to the endpoint's site.



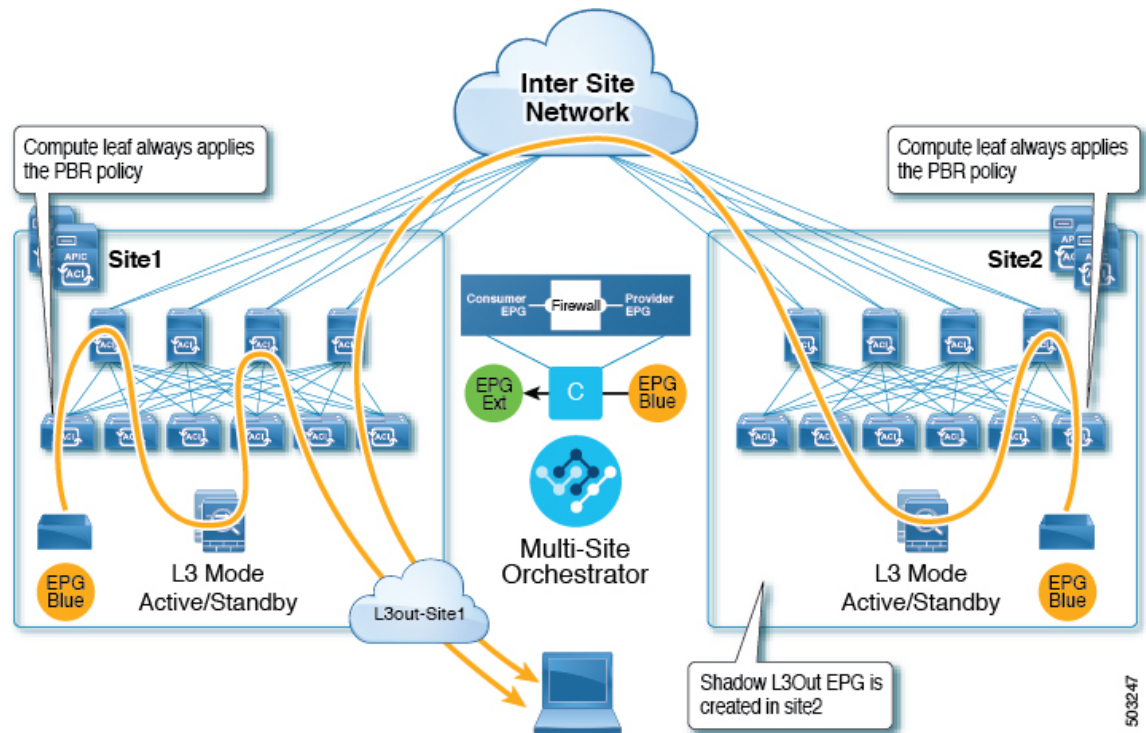
Note The same flow applies in cases when the external EPG is stretched and each site has its own L3Out, but the L3Out in the site where the traffic is originating or is destined to is down.

Figure 19: Inbound Traffic



503244

Figure 20: Outbound Traffic



Site-Local EPG

This use case illustrates a site-local application EPG that will use the L3Out in the other site for North-South traffic. Like in the previous example, all traffic will use the EPG's site-local service graph device.



Note The same flow applies in cases where the external EPG is stretched and each site has its own L3Out, but the EPG's local L3Out is down.

Figure 21: Inbound Traffic

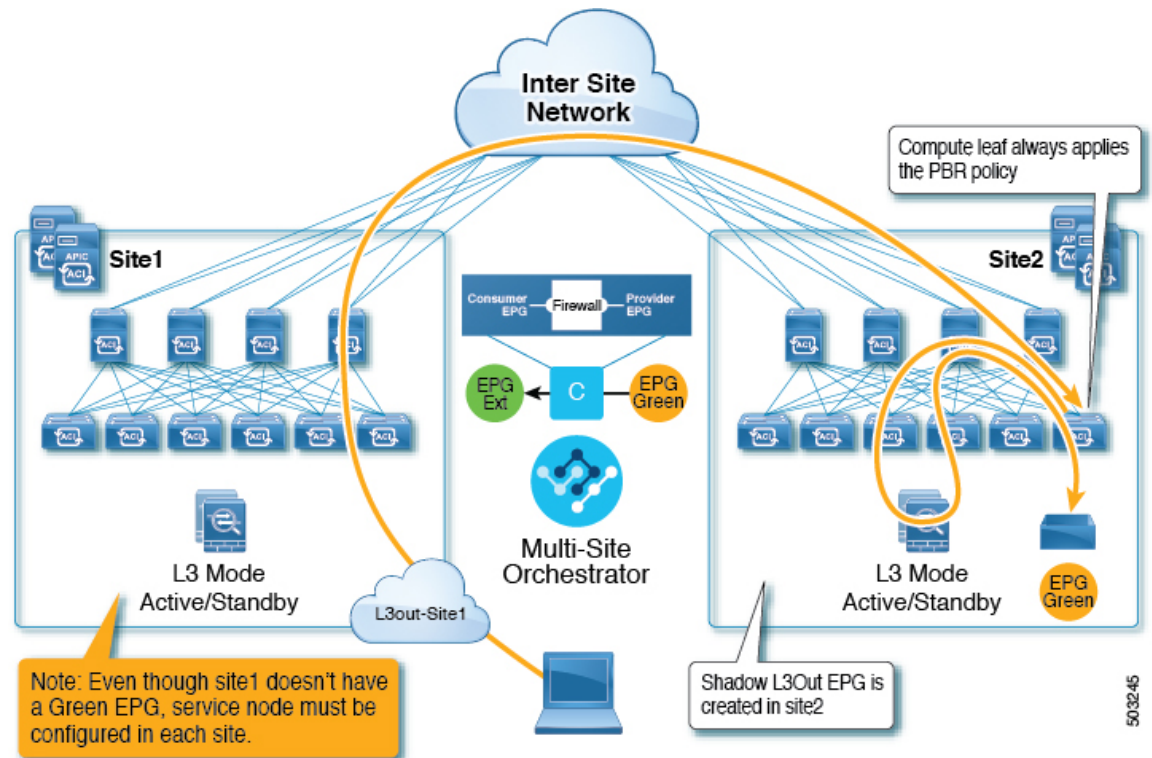
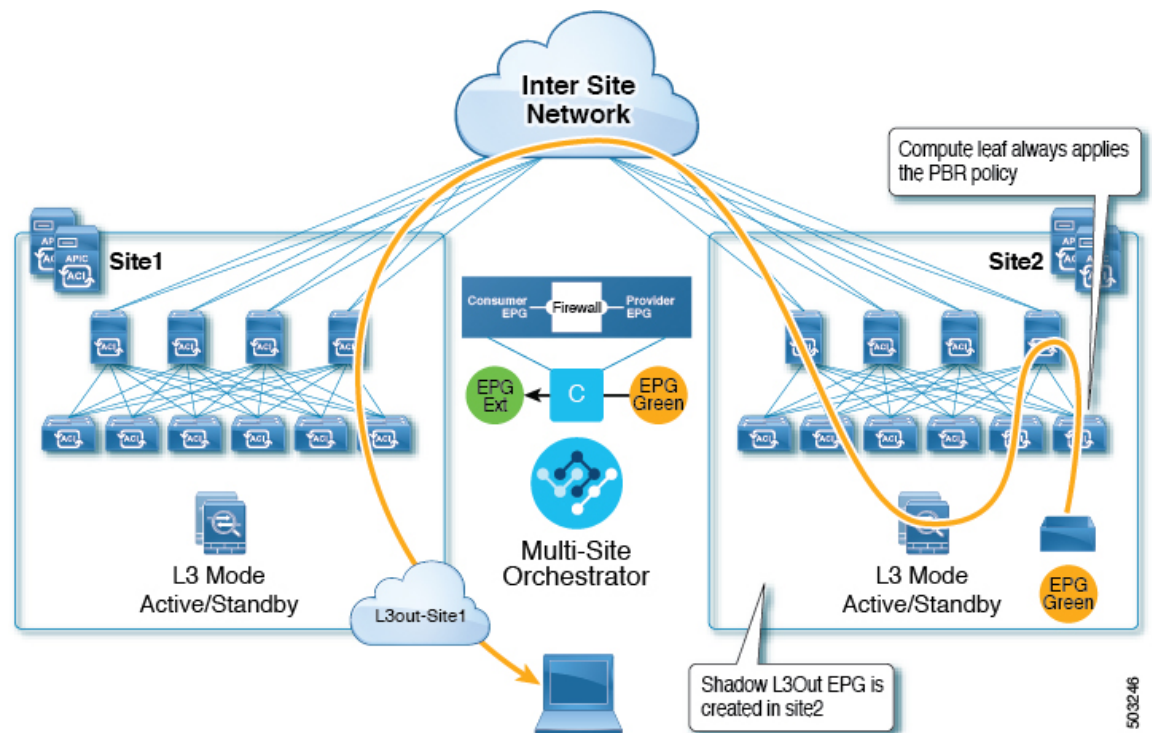


Figure 22: Outbound Traffic



Guidelines and Limitations

When configuring an Intersite L3Out with PBR, the following restrictions apply:

- When configuring MultiSite Policy-Based Redirect (PBR) between Endpoint Groups (EPGs), the following features are not supported for specific IP endpoints or host prefixes (/32 for IPv4 and /128 for IPv6):
 - Static Route on a Bridge Domain (NH Reachability)
 - Microsoft Network Load Balancing
 - Anycast MAC
- For intersite L3Out without PBR use cases, see [Intersite L3Out, on page 205](#)
- For intersite L3Out with PBR, the following use cases are supported:
 - Inter-VRF intersite L3Out with the application EPG as the `consumer`.
For inter-VRF contracts, the L3Out must be the `provider`.
 - Intra-VRF intersite L3Out with the application EPG as either the `provider` or the `consumer`
 - Intersite transit routing (L3Out-to-L3Out) with PBR is not supported.
- The above use cases are supported for sites running Cisco APIC, Release 4.2(5) or Release 5.1(x). They are not supported for sites running Cisco APIC, Release 5.0(x).
- In all supported cases, the application EPG can be stretched or not stretched.
- Service graph devices must be defined in each site, including the sites that don't have an application EPG that has a PBR contract with an intersite L3Out external EPG.
- Both one-arm and two-arm deployment models are supported.

In one-arm deployment, both the inside and outside interfaces of the service graph are connected to the same bridge domain. In two-arm deployments, the service graph interfaces are connected to separate BDs.
- When configuring a load balancer with PBR, the load balancer and the real servers for the virtual IP (VIP) must be in the same site. If PBR is disabled, the load balancer and the real servers can be in different sites.
- When configuring PBR, destination can be L1, L2, or L3.

Configuring APIC Sites

Configuring External TEP Pool

Intersite L3Out requires an external TEP address for the border leaf switches in each pod. If you already have an external TEP pool configured, for example for another feature such as Remote Leaf, the same pool can be used. The existing TEP pool will be inherited by the Nexus Dashboard Orchestrator and shown in the GUI as part of the infra configuration. Otherwise, you can add a TEP pool in the GUI, as described in this section.



Note Every pod must be assigned a unique TEP pool and it must not overlap with any other TEP pool in the fabric

Procedure

-
- Step 1** Log in to your Nexus Dashboard Orchestrator.
 - Step 2** In the left navigation menu, select **Infrastructure > Site Connectivity**.
 - Step 3** In the top right of the main pane, click **Configure**.
 - Step 4** In the left sidebar, select the site you want to configure.
 - Step 5** In the main window, click a pod in the site.
 - Step 6** In the right sidebar, click **+Add TEP Pool**.
 - Step 7** In the **Add TEP Pool** window, specify the external TEP pool you want to configure for that site.

Note

You must ensure that the TEP pool you are adding does not overlap with any other TEP pools or fabric addresses.

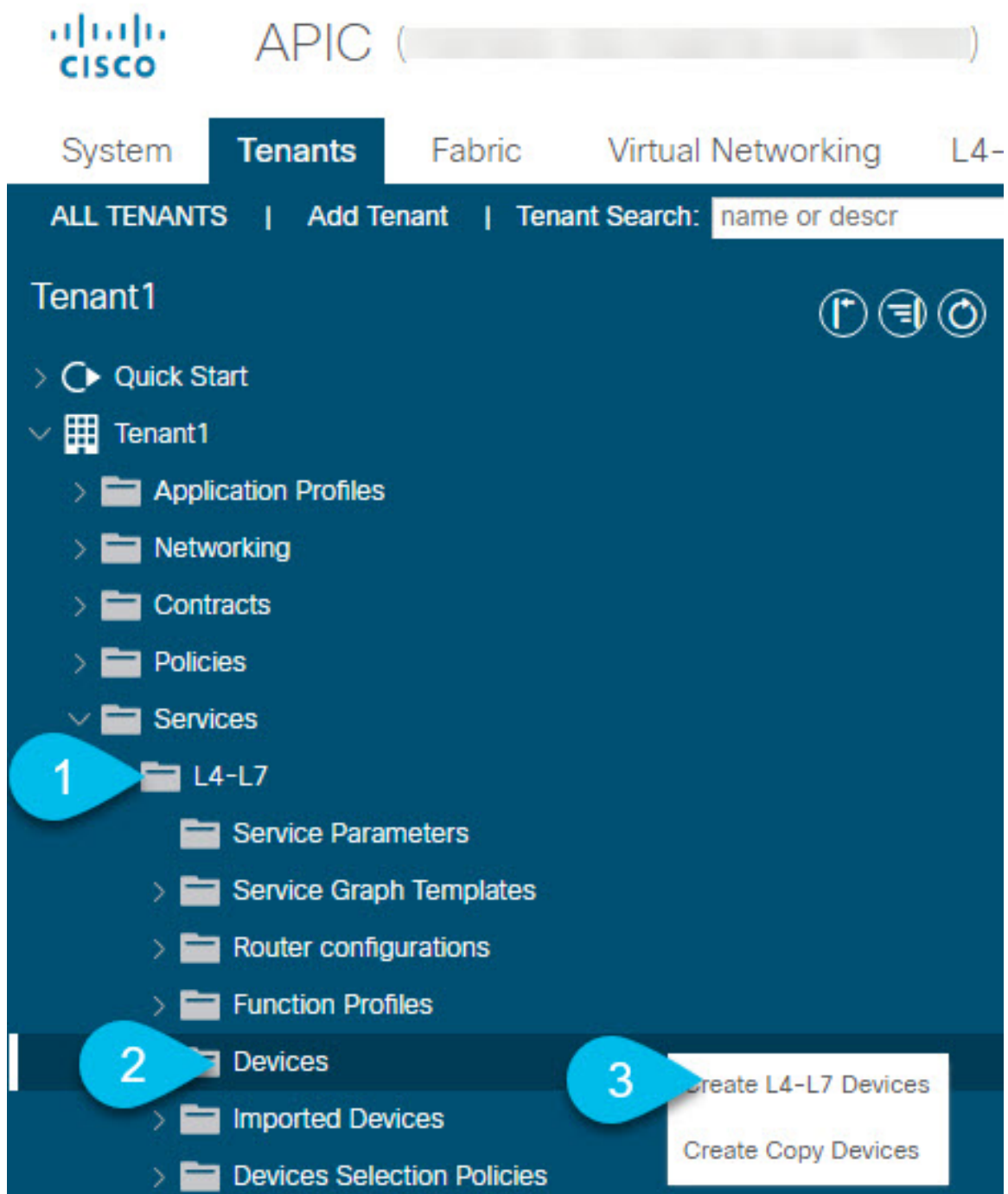
- Step 8** Repeat the process for each site and pod where you plan to use intersite L3Outs.
-

Creating and Configuring L4-L7 Devices and PBR Policies

You must create the service graph devices and define the PBR policies directly in each site's APIC.

Procedure

-
- Step 1** Log in to your Cisco APIC.
 - Step 2** In the top menu bar, click **Tenants**, then select the tenant where you want to create the device.
 - Step 3** Create an L4-L7 device.



- In the left sidebar, expand <tenant-name> > **Services** > **L4-L7** category.
- Right-click **Devices** category.
- Choose **Create L4-L7 Devices**.

The **Create L4-L7 Devices** configuration dialog opens.

Step 4 Configure the L4-L7 device.

The following image shows a sample device configuration. Your configuration settings will depend on the type and purpose of the device.

Create L4-L7 Devices

STEP 1 > General

General

Managed: ☐

Name: Site1-FW

Service Type: Firewall

Device Type: CLOUD PHYSICAL VIRTUAL

VMM Domain: S1-VMM

Trunking Port: ☐

VM Instantiation Policy: select an option

Promiscuous Mode: ☐

Context Aware: Multiple Single

Function Type: GoThrough GoTo

Devices

The device mode can be single, HA or cluster. Create only one device for single, two for HA and at least 3 for cluster.

Name	VM Name	vCenter Name	Interfaces
ASAv1	MSO-SG-WP-ASAv1	vc1	g0/0 g0/1
ASAv2	MSO-SG-WP-ASAv2	vc1	g0/0 g0/1

Cluster

Cluster Interfaces:

Name	Concrete Interfaces
FW-external	ASAv1/g0/0,ASAv2/g0/0
FW-internal	ASAv1/g0/1,ASAv2/g0/1

Previous Cancel Finish

Step 5 Create a PBR policy.

- In the left sidebar, expand <tenant-name> > **Policies** > **Protocol** category.
- Right-click **L4-L7 Policy-Based Redirect** category.
- Choose **Create L4-L7 Policy-Based Redirect**.

The **Create L4-L7 Policy-Based Redirect** configuration dialog opens.

Step 6 Configure the PBR policy.

The following image shows a sample PBR policy configuration with destination IP and MAC added.

Your configuration settings will depend on the type and purpose of the device and policy you create. For example, you can configure additional options such as IP-SLA, hashing algorithm, resilient hashing, and so on in the PBR policy.

Create L4-L7 Policy-Based Redirect

Name: FW-external

Description: optional

Destination Type:

L1

L2

L3

IP SLA Monitoring Policy: select an option

Enable Pod ID Aware Redirection: ☐

Hashing Algorithm:

dip

sip

sip-dip-prototype

Enable Anycast: ☐

Resilient Hashing Enabled: ☐

L3 Destinations:

IP	Destination Name	MAC	Redirect Health Group	Additional IPv4/IPv6	Description	Oper Status
192....		00:50:56:95:...				Ena...

Cancel

Submit

Step 7 Repeat the previous steps to create the required devices and PBR policies in the other site.

Creating Templates

When creating the schema and template, we recommend separating the templates in the following way:

- A single shared template that will contain all the objects that are stretched between all sites.
- One template per site that will contain the objects you will deploy to that site only.

In this example, we will work with two sites, so we will create a total of three templates: one for each site, plus one stretched.

Before you begin

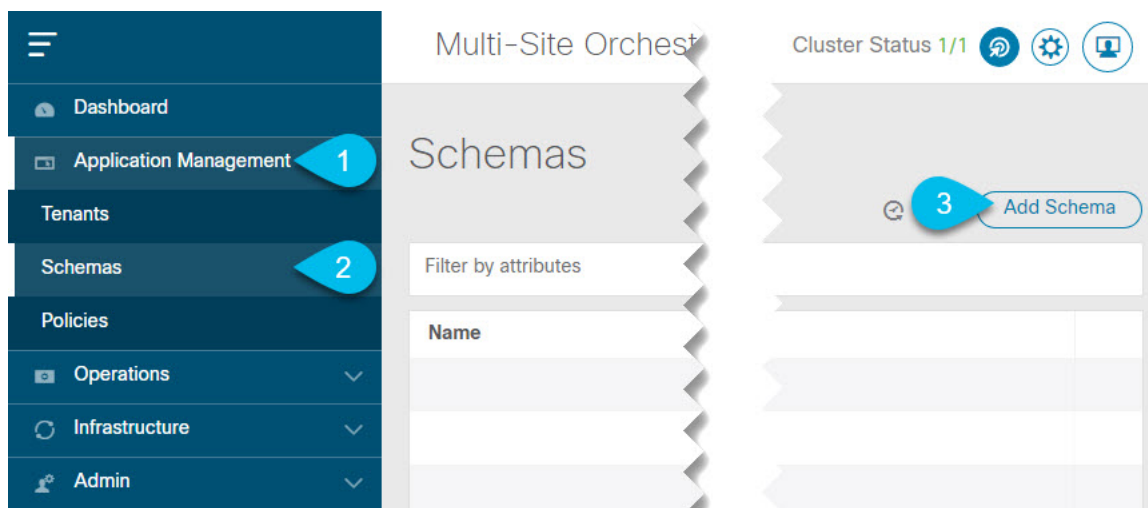
You must have:

- Reviewed the [Guidelines and Limitations](#), on page 230 and completed any prerequisites listed there.
- Finished configuring the individual APIC sites as described in [Configuring External TEP Pool](#), on page 207 and [Creating and Configuring L4-L7 Devices and PBR Policies](#), on page 231.

Procedure

Step 1 Log in to your Cisco Nexus Dashboard Orchestrator GUI.

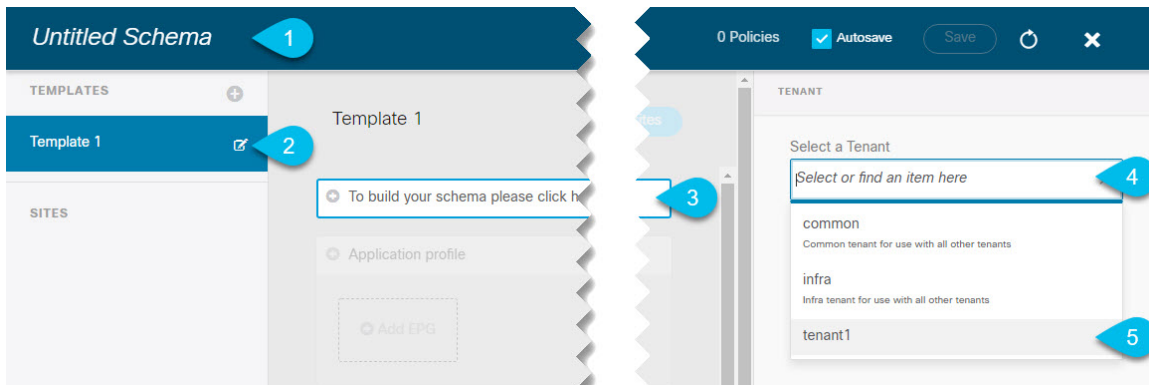
Step 2 Create a new Schema.



- In the left navigation sidebar, expand the **Application Management** category.
- Choose **Schemas**.
- Click **Add Schema** to create a new schema.

The **Edit Schema** window will open.

Step 3 Name the Schema and pick the Tenant.



- a) Replace **Untitled Schema** with the name for your schema.
Simply click on the `Untitled Schema` name to edit it.
- b) Rename the template.
In the left sidebar, mouse over the template and click the **Edit** icon.
For example, `template-stretched`.
- c) In the main pane, click **To build your schema please click here to select a tenant**.
- d) In the right sidebar, click the **Select a Tenant** dropdown.
- e) Select the tenant.

Step 4 Create any additional templates.

In the left sidebar, click the plus (+) icon next to **Templates** to add the site-specific templates. Then follow the same instructions described in the previous steps to name the templates and pick the tenant.

For example, `template-site1` and `template-site2`.

Configuring Service Graph

You must have:

- Created the L4-L7 devices directly in each site's APIC, as described in [Creating and Configuring L4-L7 Devices and PBR Policies, on page 231](#).
- Created the templates where you will create these objects, as described in [Creating Templates, on page 234](#).

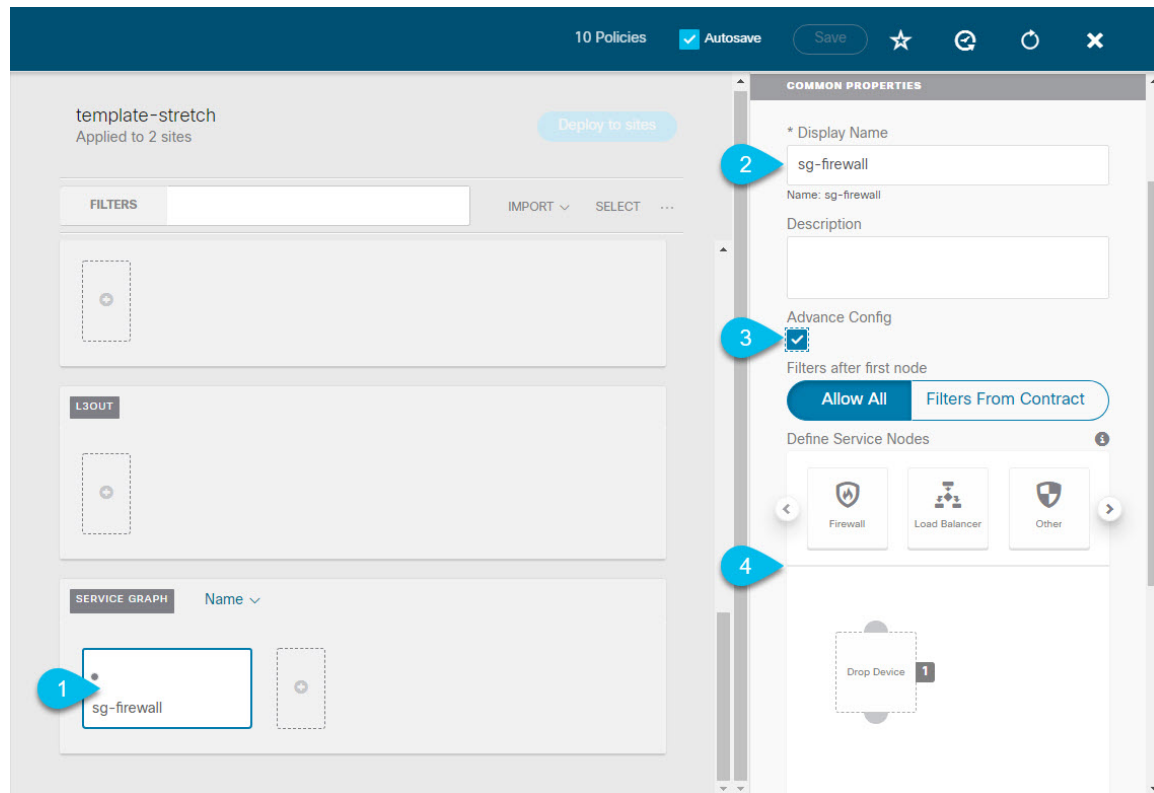
This section describes how to configure one or more devices for a service graph.

Procedure

- Step 1** Select the template where you will create the service graph.

You will create a single service graph in the `template-stretch` but configure site-local devices for it as described later in this procedure.

Step 2 Create the Service Graph.



- In the main pane, scroll down to the **Service Graph** area and click the + sign to create a new one.
- Provide the **Display Name** for the service graph.
- (Optional) Check the **Advanced Config** option.

This option allows you to configure whether traffic is restricted or not after the first service graph node. If you do not enable this option, all traffic is allowed after the first service graph node by default.

If you choose to enable the **Advanced Config**, select one of the following two options:

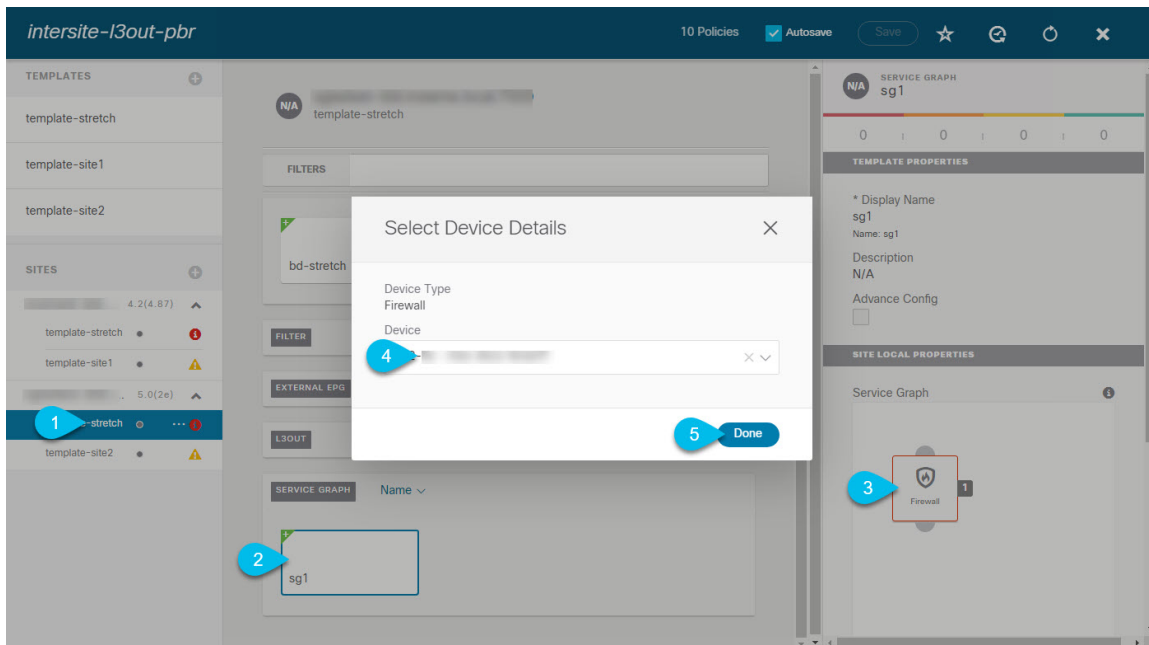
- **Allow All:** Use default (`permit-all`) filter instead of specific filter from contract subject.
This is the same behavior as with **Advanced Config** disabled.
- **Filters From Contract:** Use specific filters from contract subject.

- In the right sidebar, scroll down to the **Define Service Nodes** area and drag and drop one or more nodes into the **Drop Device** box.

Multi-Site supports up to two nodes per service graph.

Step 3 Configure service graph's site-local devices.

You must perform this step for every site that is part of the Multi-Site domain.



- From the left sidebar, select one of the sites where you will deploy this service graph.
- In the main pane, select the service graph you created.
- In the right sidebar, click on the service graph node.
- In the **Select Device Details** window, choose the device you have created in the site's APIC.

Creating Filter and Contract

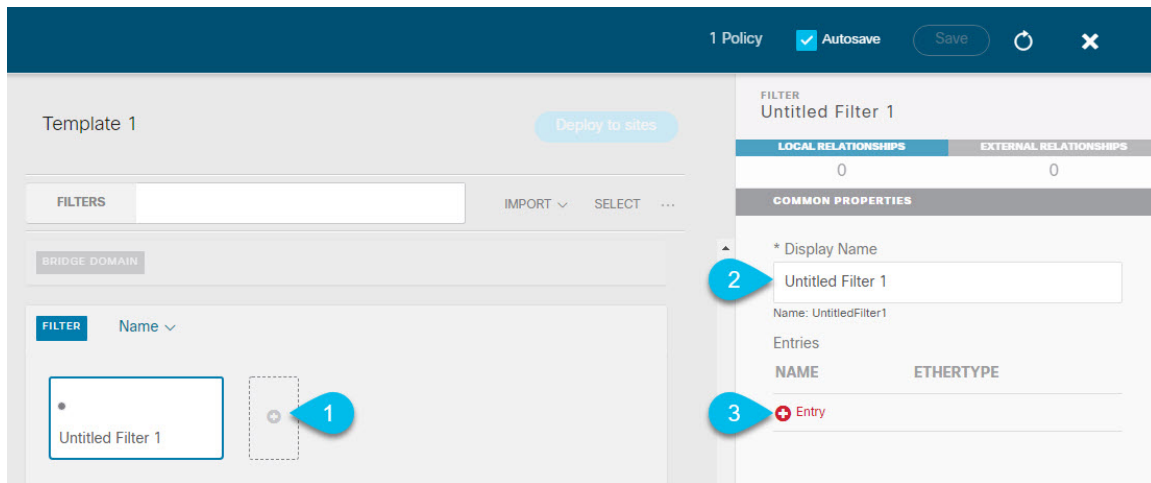
You must have:

- Created the templates where you will create these objects, as described in [Creating Templates, on page 234](#).

This section describes how to create a contract and filters that will be used for the traffic going between the application EPG and the L3Out through the service graph.

Procedure

- Step 1** Create a filter.



- In the middle pane, scroll down to the **Filter** area, then click + to create a filter.
- In the right pane, provide the **Display Name** for the filter.
- In the right pane, click + **Entry**.

Step 2 Provide the filter details.

Add Entry

×

COMMON PROPERTIES

Name

icmp

1

Description

Ether Type

ip

2

IP Protocol

icmp

Destination port range from

unspecified

Destination port range to

unspecified

3

ON-PREM PROPERTIES

☐ Match only fragments

☐ stateful

ARP flag

unspecified

×

▼

Source port range from

unspecified

▼

Source port range to

unspecified

▼

TCP session rules

unspecified

×

▼

4

Save

- Provide the **Name** for the filter.
- Choose the **Ether Type** and **IP Protocol**.

240

Cisco Nexus Dashboard Orchestrator Configuration Guide for ACI Fabrics, Release 4.0(x)

For example, `ip` and `icmp`.

- c) Leave other properties unspecified.
- d) Click **Save** to save the filter.

Step 3 Create a contract

The screenshot shows the Cisco Nexus Dashboard Orchestrator interface. The top bar indicates 14 Policies, Autosave is enabled, and there are buttons for Save, Star, Refresh, and Close. The main area is titled 'template-stretch' and 'Applied to 2 sites'. Below this, there are sections for FILTERS, Application Profile, CONTRACT, VRF, BRIDGE DOMAIN, and FILTER. The CONTRACT section is highlighted with a blue box and a red dot, with a red callout '1' pointing to it. The right pane shows the COMMON PROPERTIES section with fields for Display Name (c1), Scope (vrf), Apply both directions (checked), Filter Chain (icmp), and Service Graph (sg-firewall). A red callout '2' points to the Display Name field, '3' to the Scope dropdown, '4' to the Apply both directions toggle, '5' to the Filter Chain dropdown, and '6' to the Service Graph dropdown. Below these fields is a diagram showing a Firewall node connected to Consumer EPG and Provider EPG, with a red callout '7' pointing to the Firewall node. A red note at the bottom of the right pane states: 'Consumer and Provider connector configurations are required for service nodes. Click on node to select connector settings.'

- a) In the middle pane, scroll down to the **Contract** area and click + to create a contract.
- b) In the right pane, provide the **Display Name** for the contract
- c) From the **Scope** dropdown menu, select the scope of the contract.

If your application EPG and L3Out are in the same VRF, choose `vrf`; otherwise, if you are configure inter-VRF use case, select `tenant`.

- d) Ensure that **Apply both directions** is enabled.

This allows you to use the same filter to apply for both consumer-to-provider and provider-to-consumer directions.

- e) In the right pane, scroll down to the **Filter Chain** area and click + **Filter** to add a filter to the contract.

In the **Add Filter Chain** window that opens, select the filter you added in previous section from the **Name** dropdown menu.

If you disabled the `Apply both directions` option, repeat this stem for the other filter chain.

- f) From the **Service Graph** dropdown, select the service graph you created in previous section.
- g) Click the service graph node to configure its connectors.

Step 4 Select bridge domains for the service graph nodes' connectors.

Configure Firewall

Consumer EPG — Firewall — Provider EPG

Consumer Connector

* Bridge Domain 1

1

Provider Connector

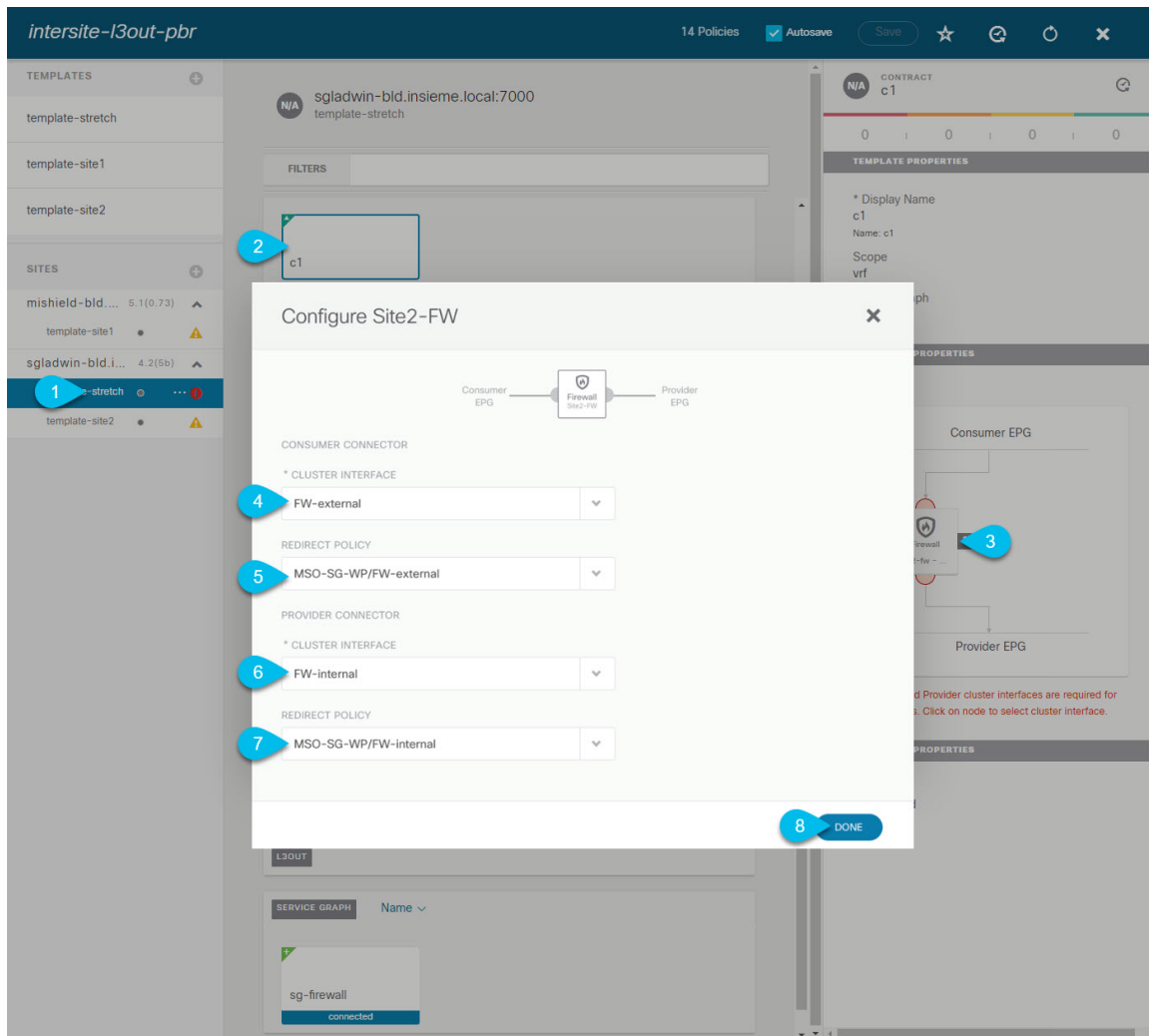
* Bridge Domain 2

2

3 Done

- a) Provide the **Consumer Connector** bridge domain.
- b) Provide the **Provider Connector** bridge domain.
- c) Click **Done** to save.

Step 5 Configure the contract's site-local properties.



- In the left sidebar, select the template under a site to which it is assigned.
- In the main pane, select the contract.
- In the right sidebar, click a service graph node.
- Select the **Cluster Interface** for the **Consumer Connector**.
- Select the **Redirect Policy** for the **Consumer Connector**.
- Select the **Cluster Interface** for the **Provider Connector**.
- Select the **Redirect Policy** for the **Provider Connector**.
- Click **Done** to save the changes.
- Repeat this step for every site.

Creating Application EPG

Creating VRF and Bridge Domain for Application EPG

This section describes how to create the VRF and bridge domain (BD) for your application EPG.

Before you begin

You must have:

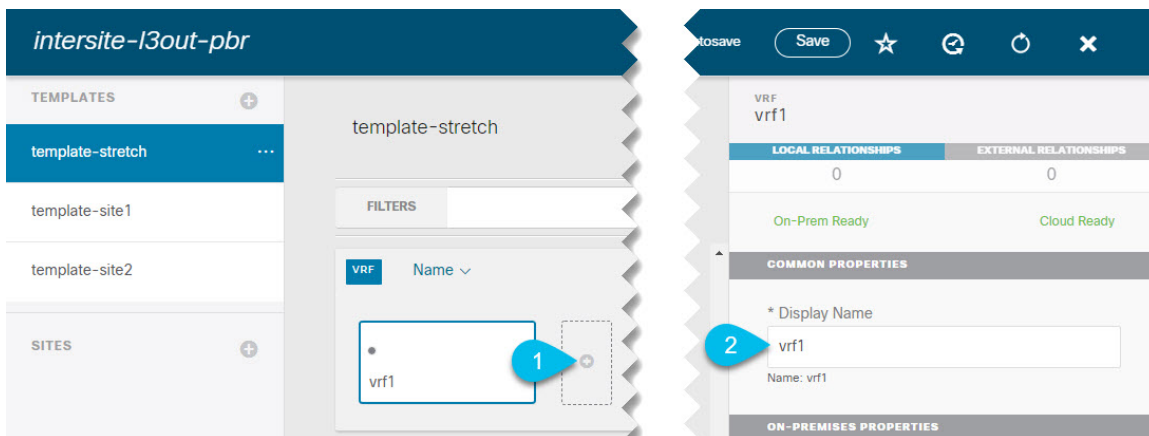
- Created the templates where you will create these objects, as described in [Creating Templates, on page 234](#).

Procedure

Step 1 Select the template where you will create the VRF and BD.

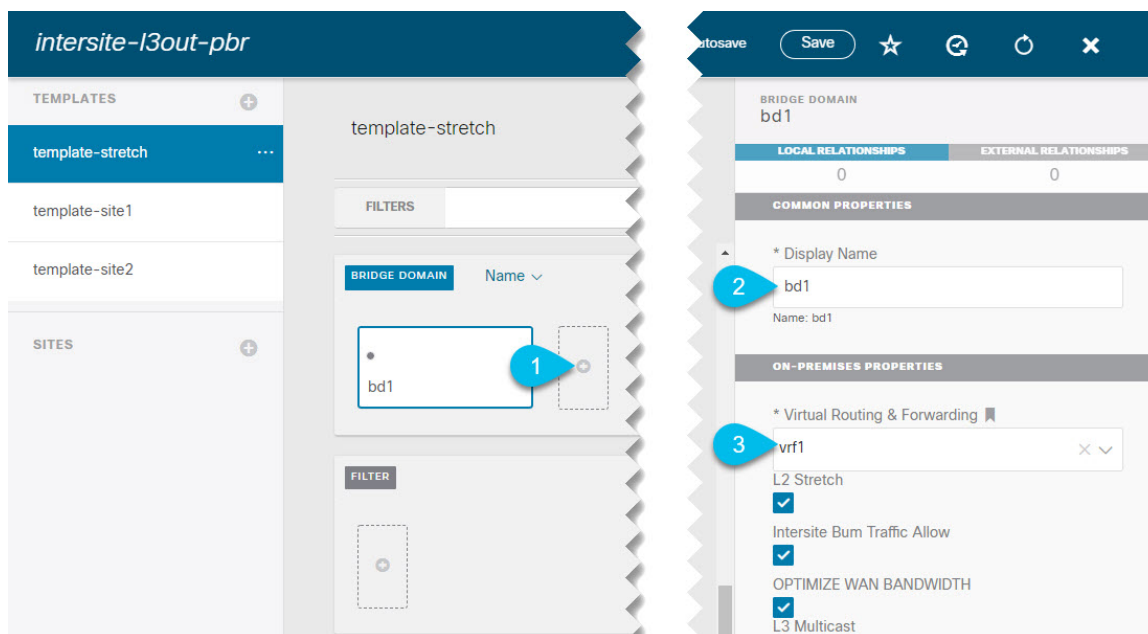
If you are planning to stretch the VRF and BD, select the `template-stretch` template. Otherwise, choose one of the site-specific templates.

Step 2 Create VRF.



- In the main pane's **VRF** area, click the plus (+) sign to add a VRF.
- In the right sidebar, provide the **Display Name** for the VRF.
- Specify other VRF settings as appropriate for your deployment.

Step 3 Create BD.



- In the main pane's **BD** area, click the plus (+) sign to add a BD.
- In the right sidebar, provide the **Display Name** for the BD.
- From the **Virtual Routing & Forwarding** dropdown, select the VRF you created in the previous step.
- Specify other BD settings as appropriate for your deployment.

Creating Application Profile and EPG

This section describes how to create the application EPG you will later configure to use the intersite L3Out with Service Graph.

Before you begin

You must have:

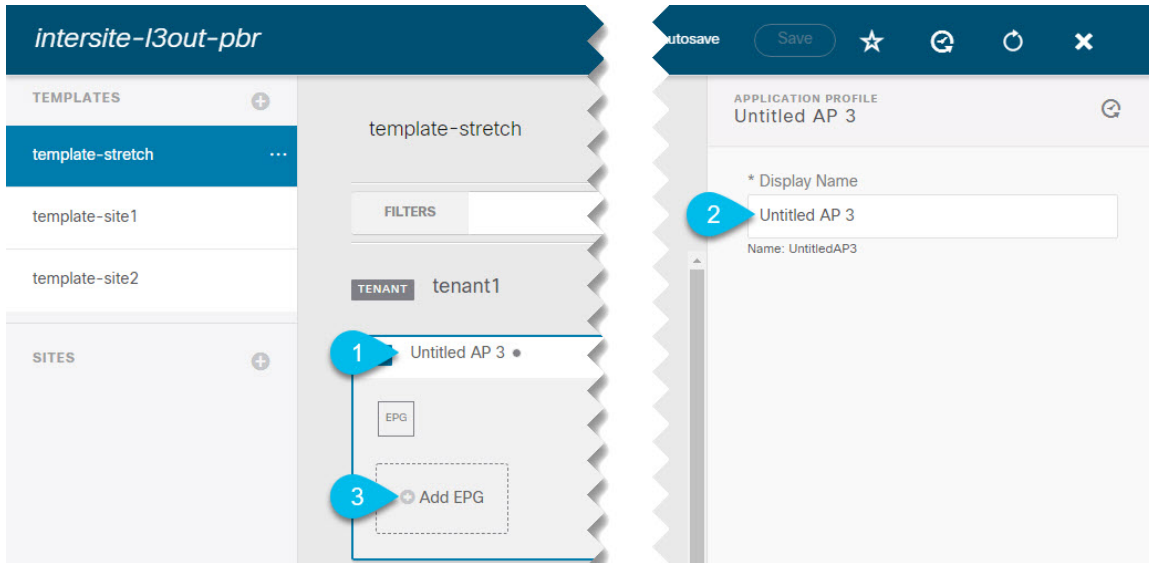
- Created the templates where you will create these objects, as described in [Creating Templates, on page 234](#).
- Created the contract you plan to use for communication between the application EPG and the external EPG, as described in [Creating Filter and Contract, on page 238](#).
- Created the VRF and BD you plan to use for the EPG, as described in [Creating VRF and Bridge Domain for Application EPG, on page 244](#).

Procedure

- Step 1** Select the template where you want to create the objects.

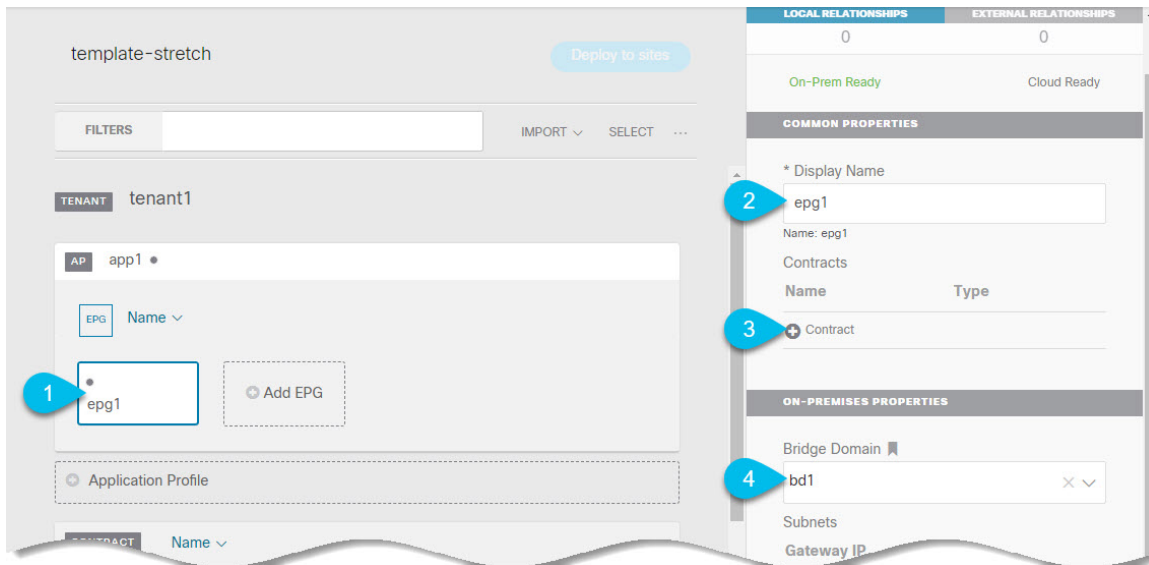
If you plan to stretch the application EPG, create it in the stretched template. If your application EPG is going to be site local, create it in the site-specific template.

Step 2 Create an application profile and EPG.



- In the main pane, click + **Application profile**.
- In the right sidebar, provide the **Display Name** for the profile.
- In the main pane, click +**Add EPG**.

Step 3 Configure the EPG.



- In the main pane, select the application EPG.
- In the right sidebar, provide the **Display Name** for the EPG.
- Click +**Contract** and select the contract.

Select the contract you have created for the EPG communication and set its type.

If you are using the same VRF for your application EPG and the L3Out external EPG, you can choose either one to be the `consumer` or the `provider`. However, if they are in different VRFs, you must select `consumer` for the application EPG's contract type.

- d) From the **Bridge Domain** dropdown, select the BD.
- e) Specify other EPG settings as appropriate for your deployment.

Creating L3Out External EPG

Creating or Importing Intersite L3Out and VRF

This section describes how to create an L3Out and associate it to a VRF in the Nexus Dashboard Orchestrator (NDO) GUI, which will then be pushed out to the APIC site, or import an existing L3Out from one of your APIC sites. You will then associate this L3Out with an external EPG and use that external EPG to configure specific intersite L3Out use cases.



Note The VRF you assign to the L3Out can be in any template or schema, but it must be in the same tenant as the L3Out.

Before you begin

You must have:

- Created the templates where you will create these objects, as described in [Creating Templates, on page 234](#).

Procedure

Step 1 Log in to your Nexus Dashboard Orchestrator.

Step 2 From the left navigation pane, select **Application Management > Schemas**.

Step 3 Select the schema and then the template where you want to create or import the VRF and L3Out.

If you create the L3Out in a template that is associated to multiple sites, the L3Out will be created on all of those sites. If you create the L3Out in a template that is associated with a single site, the L3Out will be created in that site only.

Step 4 Create a new VRF and L3Out.

If you want to import an existing L3Out, skip this step.

Note

While you can create the L3Out object in the NDO and push it out to the APIC, the physical configuration of the L3Out must be done in the APIC.

- a) Scroll down to the **VRF** area and click the + icon to add a new VRF.

In the right sidebar, provide the name for the VRF, for example `vrf-l3out`

- b) Scroll down to the **L3Out** area and click the + icon to add a new L3Out.

In the right sidebar, provide the required information.

- c) Provide the name for the L3Out, for example `l3out-intersite`.
 d) From the **Virtual Routing & Forwarding** dropdown, select the VRF you created in the previous step.

Step 5 Import an existing L3Out.

If you created a new L3Out in previous step, skip this step.

At the top of the main template view, click **Import**, then select the site from which you want to import.

Import from Site7

POLICY TYPE	SELECT TO IMPORT	INCLUDE RELATIONS
APPLICATION PROFILE 0 out of 1	☒ 1 VRF	☒
EPG 0 out of 7		
EXTERNAL EPG 0 out of 1		
CONTRACT 0 out of 5		
FILTER 0 out of 2		
VRF 1 out of 3		
BD 0 out of 15		
L3OUT 1 out of 1		

Import

- a) In the import window's **Policy Type** menu, select **L3Out**.
 b) Check the L3Out you want to import.
 c) (Optional) If you want to import all objects associated with the L3Out, enable the **Include Relations** knob.

- d) Click **Import**.
-

Configuring External EPG

This section describes how to create an external EPG that will be associated to the intersite L3Out. You can then use this external EPG and contracts to configure specific use cases for endpoints in one site to use an L3Out in another site.

Before you begin

You must have:

- Created the templates where you will create these objects, as described in [Creating Templates, on page 234](#).
- Created or imported the L3Out and VRF as described in [Creating or Importing Intersite L3Out and VRF, on page 247](#).

Procedure

Step 1 Select the template where you want to create the external EPG.

If you create the external EPG in a template that is associated to multiple sites, the external EPG will be created on all of those sites. If you create the external EPG in a template that is associated with a single site, the external EPG will be created in that site only.

Step 2 Scroll down to the **External EPG** area and click the + icon to add an external EPG.

In the right sidebar, provide the required information.

- a) Provide the name for the external EPG, for example `extEpg`.
- b) From the **Virtual Routing & Forwarding** dropdown, select the VRF you created and used for the L3Out.
- c) Click **+Contract** and select the contract.

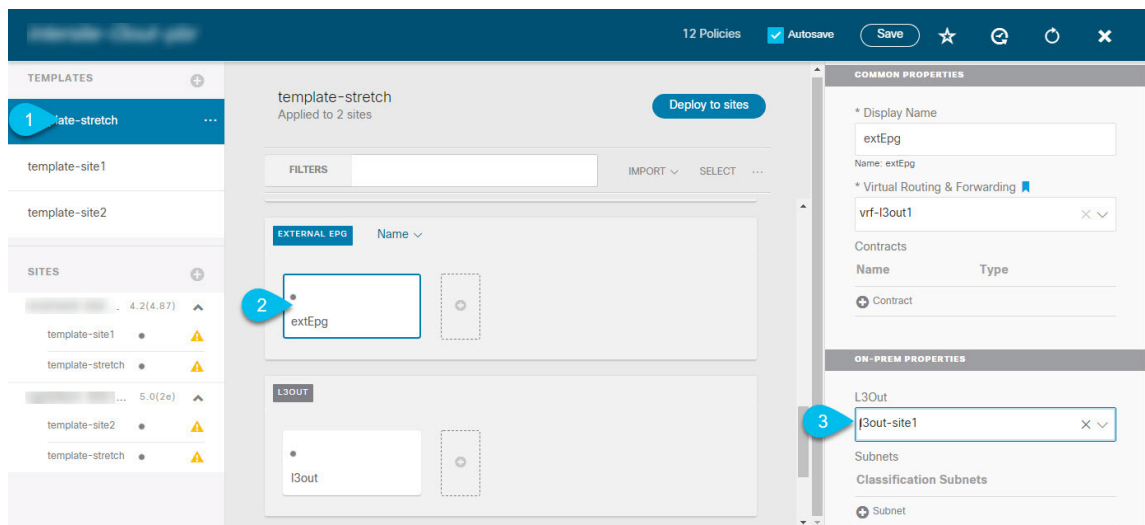
Select the contract you have created for the EPG communication and set its type.

If you are using the same VRF for your application EPG and the L3Out external EPG, you can choose either one to be the `consumer` or the `provider`. However, if they are in different VRFs, you must select `provider` for the external EPG's contract type.

Step 3 If you want to assign the L3Out at the template level...

You can choose to configure the L3Out for the external EPG at the template level, in which case, you will not be able to set the L3Outs at the site-local level.

Configuring External EPG

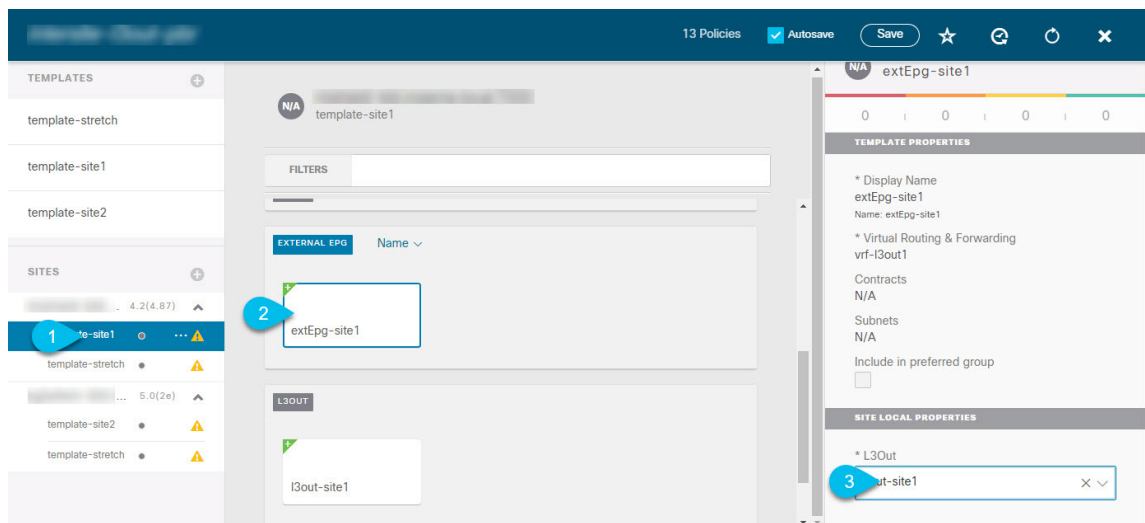


- In the left sidebar of the schema view, select the template where the external EPG is located
- Scroll down to the **External EPG** area and select the external EPG.
- In the right sidebar, scroll down to the **L3Out** dropdown and choose the intersite L3Out you created.

Step 4

If you want to assign the L3Out at the site local level...

Alternatively, you can choose to associate an L3Out with the external EPG at the site-local level.



- In the left sidebar of the schema view, select the site where the external EPG is deployed.
- Scroll down to the **External EPG** area and select the external EPG.
- In the right sidebar, scroll down to the **L3Out** dropdown and choose the intersite L3Out you created.

In this case, both the APIC-managed and the Orchestrator-managed L3Outs will be available for selection. You can select either the L3Out you have created in the previous section specifically for this or pick an L3Out that exists in the site's APIC.



CHAPTER 21

Layer 3 Multicast

- [Layer 3 Multicast, on page 251](#)
- [Layer 3 Multicast Routing, on page 252](#)
- [Rendezvous Points, on page 252](#)
- [Multicast Filtering, on page 253](#)
- [Layer 3 Multicast Guidelines and Limitations, on page 254](#)
- [Creating Multicast Route Map Policy, on page 255](#)
- [Enabling Any-Source Multicast \(ASM\) Multicast, on page 257](#)
- [Enabling Source-Specific Multicast \(SSM\), on page 258](#)

Layer 3 Multicast

Cisco Multi-Site Layer 3 multicast is enabled or disabled at three levels, the VRF, the bridge domain (BD), and any EPGs that have multicast sources present.

At the top level, multicast routing must be enabled on the VRF that has any multicast-enabled BDs. On a multicast-enabled VRF, there can be a combination of multicast-enabled BDs and BDs where multicast routing is disabled. Enabling multicast routing on a VRF from the Cisco Nexus Dashboard Orchestrator GUI enables it on the APIC sites where the VRF is stretched.

Once a VRF is enabled for multicast, the individual BDs under that VRF can be enabled for multicast routing. Configuring Layer 3 multicast on a BD enables protocol independent multicast (PIM) routing on that BD. By default, PIM is disabled in all BDs.

If a source belonging to a specific site-local EPG sends multicast traffic to a remote site, the Nexus Dashboard Orchestrator must create a shadow EPG and program the corresponding subnet route(s) on the remote site for the source EPG. In order to limit the configuration changes applied to the remote Top-of-Rack (TOR) switches, you are required to explicitly enable Layer 3 multicast on the local EPGs that have multicast sources present, so that only the configuration necessary for those EPGs is pushed to the remote sites. EPGs with multicast receivers do not require enabling Layer 3 multicast.

Multi-Site supports all of the following Layer 3 multicast source and receiver combinations:

- Multicast sources and receivers inside ACI fabric
- Multicast sources and receivers outside ACI fabric
- Multicast sources inside ACI fabric with external receivers
- Multicast receivers inside ACI fabric with external sources

Layer 3 Multicast Routing

The following is a high level overview of the Layer 3 multicast routing across sites:

- When the multicast source is attached to the ACI fabric as an endpoint (EP) at one site and starts streaming a multicast flow, the specific site's spine switch that is elected as designated forwarder for the source VRF will forward the multicast traffic to all the remote sites where the source's VRF is stretched using Head End Replication (HREP). If there are no receivers in a specific remote site for that specific group, the traffic gets dropped on the receiving spine node. If there is at least a receiver, the traffic is forwarded into the site and reaches all the leaf nodes where the VRF is deployed and at that point is pruned/forwarded based on the group membership information.
- Prior to Cisco ACI Release 5.0(1), the multicast routing solution required external multicast routers to be the Rendezvous Points (RPs) for PIM-SM any-source multicast (ASM) deployments. Each site must point to the same RP address for a given stretched VRF. The RP must be reachable on each site via the site's local L3Out.
- When the source is outside and the receiver is within a fabric, the receiver will pull traffic via site's local L3Out as PIM joins toward RP and source are always sent via site local L3Out.
- Receivers in each site are expected to draw traffic from an external source via the site's local L3Out. As such, traffic received on the L3Out on one site should not be sent to other sites. This is achieved on the spine by pruning multicast traffic from being replicated into HREP tunnels.

In order to be able to do so, all multicast traffic originated from an external source and received on a local L3Out is remarked with a special DSCP value in the outer VXLAN header. The spines can hence match that specific DSCP value preventing the traffic from being replicated toward the remote sites.

- Traffic originated from a source connected to a site can be sent toward external receivers via a local L3Out or via L3Outs deployed in remote sites. The specific L3Out that is used for this solely depends on which site received the PIM Join for that specific multicast group from the external network.
- When multicast is enabled on a BD and an EPG on the Nexus Dashboard Orchestrator, all of the BD's subnets are programmed in the routing tables of all the leaf switches, including the border leaf nodes (BLs). This enables receivers attached to the leaf switches to determine the reachability of the multicast source in cases where the source BD is not present on the leaf switches. The subnet is advertised to the external network if there is a proper policy configured on the BLs. The /32 host routes are advertised if host-based routing is configured on the BD.

For additional information about multicast routing, see the [IP Multicast](#) section of the *Cisco APIC Layer 3 Networking Configuration Guide*.

Rendezvous Points

Multicast traffic sources send packets to a multicast address group, with anyone joining that group able to receive the packets. Receivers that want to receive traffic from one or more groups can request to join the group, typically using Internet Group Management Protocol (IGMP). Whenever a receiver joins a group, a multicast distribution tree is created for that group. A Rendezvous Point (RP) is a router in a PIM-SM multicast domain that acts as a shared root for a multicast shared tree.

The typical way to provide a redundant RP function in a network consists in deploying a functionality called Anycast RP, which allows two or more RPs in the network to share the same anycast IP address. This provides

for redundancy and load balancing. Should one RP device fail, the other RP can take over without service interruption. Multicast routers can also join the multicast shared tree by connecting to any of the anycast RPs in the network, with PIM `join` requests being forwarded to the closest RP.

Two types of RP configurations are supported from Nexus Dashboard Orchestrator:

- **Static RP**—If your RP is outside the ACI fabric.
- **Fabric RP**—If the border leaf switches in the ACI fabric will function as the anycast RPs.

Any number of routers can be configured to work as RPs and they can be configured to cover different group ranges. When defining the RP inside the ACI fabric, you can configure which groups the RP covers by creating a route-map policy that contains the list of groups and attaching this policy to the RP when adding it to the VRF. Creating a route map is described in [Creating Multicast Route Map Policy, on page 255](#), while VRF configuration is described in [Enabling Any-Source Multicast \(ASM\) Multicast, on page 257](#).

Both static and fabric RPs require PIM-enabled border leaf switches in the VRF where multicast routing is enabled. L3Out configuration is currently configured locally from the APIC at each site including enabling PIM for the L3Out. Please refer to the [Cisco APIC Layer 3 Networking Configuration Guide](#) for details on configuration PIM on L3Outs

Multicast Filtering

Multicast filtering is a data plane filtering feature for multicast traffic available starting with Cisco APIC, Release 5.0(1) and Nexus Dashboard Orchestrator, Release 3.0(1).

Cisco APIC supports control plane configurations that can be used to control who can receive multicast feeds and from which sources. In some deployments, it may be desirable to constrain the sending and/or receiving of multicast streams at the data plane level. For example, you may need to allow multicast senders in a LAN to be able to send only to specific multicast groups or to allow receivers to receive multicast from only specific sources.

To configure multicast filtering from the Nexus Dashboard Orchestrator, you create source and destination multicast route maps, each of which contains one or more filter entries based on the multicast traffic's source IP and/or group with an action (`Permit` or `Deny`) attached to it. You then enable the filtering on a bridge domain by attaching the route maps to it.

When creating a multicast route map, you can define one or more filter entries. Some entries can be configured with a `Permit` action and other entries can be configured with a `Deny` action, all within the same route map. For each entry, you can provide a **Source IP** and a **Group IP** to define the traffic that will match the filter. You must provide at least one of these fields, but can choose to include both. If one of the fields is left blank, it will match all values.

You can enable both multicast source filtering and multicast receiver filtering on the same bridge domain. In this case one bridge domain can provide filtering for both, the source as well as the receivers.

If you do not provide a route map for a BD, the default action is to allow all multicast traffic on the bridge domain. However, if you do select a route map, the default action changes to deny any traffic not explicitly matched to a filter entry in the route map.

Source Filtering

For any multicast sources that are sending traffic on a bridge domain, you can configure a route map policy with one or more source and group IP filters defined. The traffic is then matched against every entry in the route map and one of the following actions takes place:

- If the traffic matches a filter entry with a `Permit` action in the route map, the bridge domain will allow traffic from that source to that group.
- If the traffic matches a filter entry with a `Deny` action in the route map, the bridge domain will block traffic from that source to that group.
- If the traffic does not match any entries in the route map, the default `Deny` action is applied.

Source filter is applied to the bridge domain on the First-Hop Router (FHR), represented by the ACI leaf node where the source is connected. The filter will prevent multicast from being received by receivers in different bridge domains, the same bridge domain, and external receivers.

Destination (Receiver) Filtering

Destination (receiver) filtering does not prevent receivers from joining a multicast group. The multicast traffic is instead allowed or dropped in the data plane based on the source IP and multicast group combination.

Similarly to the source filtering, when multicast traffic matches a destination filter, one of the following actions takes place:

- If the traffic matches a filter entry with a `Permit` action in the route map, the bridge domain will allow the traffic from the multicast group to the receiver.
- If the traffic matches a filter entry with a `Deny` action in the route map, the bridge domain will block the traffic from the multicast group to the receiver.
- If the traffic does not match any entries in the route map, the default `Deny` action is applied.

Destination filter is applied to the bridge domain on the Last-Hop Router (LHR), represented by the ACI leaf node where the receiver is connected, so other bridge domains can still receive the multicast traffic.

Layer 3 Multicast Guidelines and Limitations

Up to the current software release, Cisco Nexus Dashboard Orchestrator cannot be used to deploy specific multicast control plane filtering policies, such as IGMP or PIM related policies, on each site. As such you must configure any additional policies required for your use case on each APIC site individually for end-to-end solution to work. For specific information on how to configure those settings on each site, see the [Cisco APIC Layer 3 Networking Configuration Guide](#).

You must also ensure that QoS DSCP translation policies in all fabrics are configured consistently. When you create custom QoS policies in ACI fabrics, you can create a mapping between the ACI QoS Levels and the packet header DSCP values for packets ingressing or egressing the fabric. The same ACI QoS Levels must be mapped to the same DSCP values on all sites for the multicast traffic to transit between those sites. For specific information on how to configure those settings on each site, see the [Cisco APIC and QoS](#)

Multicast Filtering

The following additional guidelines apply if you enable the multicast filtering:

- Multicast filtering is supported only for IPv4.
- You can enable either the multicast source filtering, or the receiver filtering, or both on the same bridge domain.
- If you do not want to have multicast filters on a bridge domain, do not configure a source filter or destination filter route maps on that bridge domain.

By default, no route maps are associated with a bridge domain, which means that all multicast traffic is allowed. If a route map is associated with a bridge domain, only the permit entries in that route map will be allowed, while all other multicast traffic will be blocked.

If you attach an empty route map to a bridge domain, route maps assume a `deny-all` by default, so all sources and groups will be blocked on that bridge domain.

- Multicast filtering is done at the BD level and apply to all EPGs within the BD. As such you cannot configure different filtering policies for different EPGs within the same BD. If you need to apply filtering more granularly at the EPG level, you must configure the EPGs in separate BDs.
- Multicast filtering is intended to be used for Any-Source Multicast (ASM) ranges only. Source-Specific Multicast (SSM) is not supported for source filtering and is supported only for receiver filtering.
- For both, source and receiver filtering, the route map entries are matched based on the specified `order` of the entry, with lowest number matched first. This means that lower order entries will match first, even if they are not the longest match in the list, and higher order entries will not be considered.

For example, if you have the following route map for the `192.0.3.1/32` source:

Order	Source IP	Action
1	192.0.0.0/16	Permit
2	192.0.3.0/24	Deny

Even though the second entry (`192.0.3.0/24`) is a longer match as a source IP, the first entry (`192.0.0.0/16`) will be matched because of the lower order number.

Creating Multicast Route Map Policy

This section describes how to create a multicast route map policy. You may want to create a route map for one of the following reasons:

- Define a set of filters for multicast source filtering.
- Define a set of filters for multicast destination filtering.
- Define a set of group IPs for a Rendezvous Point (RP).

When configuring an RP for a VRF, if you do not provide a route map, the RP will be defined for the entire multicast group range (`224.0.0.0/4`). Alternatively, you can provide a route map with a group or group range defined to limit the RP to those groups only.

Procedure

Step 1 Log in to your Nexus Dashboard and open the Nexus Dashboard Orchestrator service.

Step 2 Create a new Tenant Policy.

- a) From the left navigation pane, choose **Application Management > Tenant Policies**.
- b) On the **Tenant Policy Templates** page, click **Add Tenant Policy Template**.
- c) In the Tenant Policies page's right properties sidebar, provide the **Name** for the tenant.
- d) From the **Select a Tenant** dropdown, choose the tenant with which you want to associate this template.

All the policies you create in this template as described in the following steps will be associated with the selected tenant deployed to it when you push the template to a specific site.

By default, the new template is empty, so you need to add one or more tenant policies as described in the following steps. Note that you don't have to create every policy available in the template – you can create a template with just a single route map policy for your multicast use case.

Step 3 Create a Route Map Policy for Multicast.

- a) From the **+Create Object** dropdown, select **Route Map Policy for Multicast**.
- b) In the right properties sidebar, provide the **Name** for the policy.
- c) (Optional) Click **Add Description** and provide a description for the policy.
- d) Click **+Add Route Map for Multicast Entries** and provide the route map information.

For each route map, you need to create one or more route map entries. Each entry is a rule that defines an action based on one or more matching criteria based on the following information:

- **Order** – Order is used to determine the order in which the rules are evaluated.
- **Group IP, Src IP, and RP IP** – You can use the same multicast route map policy UI for two different use cases—to configure a set of filters for multicast traffic or to restrict a rendezvous point configuration to a specific set of multicast groups. Depending on which use case you're configuring, you only need to fill some of the fields in this screen:
 - For multicast filtering, you can use the **Source IP** and the **Group IP** fields to define the filter. You must provide at least one of these fields, but can choose to include both. If one of the fields is left blank, it will match all values.

The Group IP range must be between 224.0.0.0 and 239.255.255.255 with a netmask between /4 and /32. You must provide the subnet mask.

The **RP IP** (Rendezvous Point IP) is not used for multicast filtering route maps, so leave this field blank.
 - For Rendezvous Point configuration, you can use the **Group IP** field to define the multicast groups for the RP.

The Group IP range must be between 224.0.0.0 and 239.255.255.255 with a netmask between /4 and /32. You must provide the subnet mask.

For Rendezvous Point configuration, the **RP IP** is configured as part of the RP configuration. If a route-map is used for group filtering it is not necessary to configure an RP IP address in the route-map. In this case, leave the **RP IP** and **Source IP** fields empty.

- **Action** – Action defines the action to perform, either **Permit** or **Deny** the traffic, if a match is found.

- e) Click the checkmark icon to save the entry.
 - f) Repeat the previous substeps to create any additional route map entries for the same policy.
 - g) Click **Save** to save the policy and return to the template page.
 - h) Repeat this step to create any additional Route Map for Multicast policies.
-

Enabling Any-Source Multicast (ASM) Multicast

The following procedure describes how to enable ASM multicast on VRF, BD, and EPG using the Nexus Dashboard Orchestrator GUI. If you want to enable SSM multicast, follow the steps in [Enabling Source-Specific Multicast \(SSM\)](#), on page 258 instead.

Before you begin

- Ensure you have read and followed the information described in [Layer 3 Multicast Guidelines and Limitations](#), on page 254.
- If you plan to enable multicast filtering, create the required multicast route maps, as described in [Creating Multicast Route Map Policy](#), on page 255.
- Note that the site-local L3Outs must have PIM enabled in the VRF when fabric RP is enabled.

This is described in Step 6 of the following procedure. Additional information about PIM configuration on an L3Out is available in the [Cisco APIC Layer 3 Networking Configuration Guide](#).

Procedure

- Step 1** Log in to your Nexus Dashboard Orchestrator.
- Step 2** From the left-hand sidebar, select the **Application Management > Schemas** view.
- Step 3** Click on the Schema you want to change.
- Step 4** Enable Layer 3 multicast on a VRF.
First, you enable Layer 3 multicast on a VRF that is stretched between sites.
 - a) Select the VRF for which you want to enable Layer 3 multicast.
 - b) In the right properties sidebar, check the **L3 Multicast** checkbox.
- Step 5** Add one or more Rendezvous Points (RP).
 - a) Select the VRF.
 - b) In the right properties sidebar, click **Add Rendezvous Points**.
 - c) With the VRF still selected, click **Add Rendezvous Points** in the right sidebar.
 - d) In the **Add Rendezvous Points** window, provide the IP address of the RP.
 - e) Choose the type of the RP.
 - **Static RP**—If your RP is outside the ACI fabric.
 - **Fabric RP**—If your RP is inside the ACI fabric.

- f) (Optional) From the **Multicast Route-Map Policy** dropdown, select the route-map policy you configured previously.

By default, the RP IP you provide applies to all multicast groups in the fabric. If you want to restrict the RP to only a specific set of multicast groups, define those groups in a route map policy and select that policy here.

Step 6 Enable PIM on the L3Out.

Both static and fabric RPs require PIM-enabled border leaf switches where multicast routing is enabled. L3Out configuration currently cannot be done from the Nexus Dashboard Orchestrator, so you must ensure that PIM is enabled directly in the site's APIC. Additional information about PIM configuration on an L3Out is available in the [Cisco APIC Layer 3 Networking Configuration Guide](#).

- Log in to your site's Cisco APIC.
- In the top menu, click **Tenants** and select the tenant that contains the L3Out.
- In the left navigation menu, select **Networking > L3Outs > <l3out-name>**.
- In the main pane, choose the **Policy** tab.
- Check the **PIM** options.

Multi-Site supports IPv4 multicast only.

Step 7 Enable Layer 3 multicast on a BD.

Once you have enabled L3 Multicast on a VRF, you can enable L3 multicast on a Bridge Domain (BD) level.

- Select the BD for which you want to enable Layer 3 multicast.
- In the right properties sidebar, check the **L3 Multicast** checkbox.

Step 8 (Optional) If you want to configure multicast filtering, provide the route-maps for source and destination filtering.

- Select the BD.
- In the right properties sidebar, select a **Route-Map Source Filter** and **Route-Map Destination Filter**.

You can choose to enable either the multicast source filtering, or the receiver filtering, or both.

Keep in mind, if you do not select a route map, the default action is to allow all multicast traffic on the bridge domain; however, if you select a route map the default action changes to deny any traffic not explicitly matched to a filter entry in the route map.

Step 9 If your multicast source is in one site and is not stretched to the other sites, enable intersite multicast source option on the EPG.

Once you have enabled L3 Multicast on the BD, you must also enable multicast on the EPGs (part of multicast-enabled BDs) where multicast sources are connected.

- Select the EPG for which you want to enable Layer 3 multicast.
- In the right-hand sidebar, check the **Intersite Multicast Source** checkbox.

Enabling Source-Specific Multicast (SSM)

The following procedure describes how to enable SSM multicast on VRF, BD, and EPG using the Nexus Dashboard Orchestrator GUI. If you want to enable ASM multicast, follow the steps in [Enabling Any-Source Multicast \(ASM\) Multicast, on page 257](#) instead.

Before you begin

- Ensure you have read and followed the information described in [Layer 3 Multicast Guidelines and Limitations, on page 254](#).
- If you plan to enable multicast filtering, create the required multicast route maps, as described in [Creating Multicast Route Map Policy, on page 255](#).
- Note that you need to configure IGMPv3 interface policy for the multicast-enabled BDs at the site-local level.

This is described in Step 8 of the following procedure. Additional information is available in the [Cisco APIC Layer 3 Networking Configuration Guide](#).

Procedure

Step 1 Log in to your Nexus Dashboard Orchestrator.

Step 2 From the left-hand sidebar, select the **Application Management > Schemas** view.

Step 3 Click on the Schema you want to change.

Step 4 Enable Layer 3 multicast on a VRF.

First, you enable Layer 3 multicast on a VRF that is stretched between sites.

- a) Select the VRF for which you want to enable Layer 3 multicast.
- b) In the right properties sidebar, check the **L3 Multicast** checkbox.

Step 5 (Optional) Configure a custom range for SSM listeners.

The default SSM range is `232.0.0.0/8`, which is automatically configured on the switches in your fabric. If you are using SSM, we recommend configuring your listeners to join groups in this range, in which case you can skip this step.

If for any reason you do not want to change your listener configuration, you can add additional SSM ranges under the VRF settings by creating a route-map with up to 4 additional ranges. Keep in mind that if you add a new range it will become an SSM range and cannot be used for ASM at the same time.

Custom SSM range configuration must be done directly in the site's APIC:

- a) Log in to your site's Cisco APIC.
- b) In the top menu, click **Tenants** and select the tenant that contains the VRF.
- c) In the left navigation menu, select **Networking > VRFs > <VRF-name> > Multicast**.
- d) In the main pane, choose the **Pattern Policy** tab.
- e) From the **Route Map** dropdown in the **Source Specific Multicast (SSM)** area, choose an existing route map or click **Create Route Map Policy for Multicast** option to create a new one.

If you select an existing route map, click the icon next to the dropdown to view the route map's details.

In the route map details window or the **Create Route Map Policy for Multicast** window that opens, click + to add an entry. Then configure the Group IP; you need to provide only the group IP address to define the new range.

Step 6 (Optional) Enable PIM on the site's L3Out.

If you connect multicast sources and/or receivers to the external network domain, you must also enable PIM on the site's L3Out. L3Out configuration currently cannot be done from the Nexus Dashboard Orchestrator, so you must ensure

that PIM is enabled directly in the site's APIC. Additional information about PIM configuration on an L3Out is available in the [Cisco APIC Layer 3 Networking Configuration Guide](#).

- a) Log in to your site's Cisco APIC.
- b) In the top menu, click **Tenants** and select the tenant that contains the L3Out.
- c) In the left navigation menu, select **Networking > L3Outs > <l3out-name>**.
- d) In the main pane, choose the **Policy** tab.
- e) Check the **PIM** options.

Multi-Site supports IPv4 multicast only.

Step 7 Enable Layer 3 multicast on a BD.

Once you have enabled L3 Multicast on a VRF, you can enable L3 multicast on a Bridge Domain (BD) level.

- a) Select the BD for which you want to enable Layer 3 multicast.
- b) In the right properties sidebar, check the **L3 Multicast** checkbox.

Step 8 Enabled IGMPv3 interface policy on the bridge domains where receivers are connected.

Because you are configuring SSM, you must also assign an IGMPv3 interface policy to the BD. By default, when PIM is enabled, IGMP is also automatically enabled on the SVI but the default version is set to IGMPv2. You must explicitly set the IGMP interface policy to IGMPv3. This must be done at the site-local level:

- a) Log in to your site's Cisco APIC.
- b) In the top menu, click **Tenants** and select the tenant that contains the BD.
- c) In the left navigation menu, select **Networking > Bridge Domains > <BD-name>**.
- d) In the main pane, choose the **Policy** tab.
- e) From the **IGMP Policy** dropdown, select the IGMP policy or click **Create IGMP Interface Policy** to create a new one.

If you select an existing policy, click the icon next to the dropdown to view the policy details.

In the policy details window or the **Create Route Map Policy for Multicast** window that opens, ensure that the **Version** field is set to `Version 3`.

Step 9 (Optional) If you want to configure multicast filtering, provide the route-maps for source and destination filtering.

- a) Select the BD.
- b) In the right properties sidebar, select a **Route-Map Source Filter** and **Route-Map Destination Filter**.

You can choose to enable either the multicast source filtering, or the receiver filtering, or both.

Keep in mind, if you do not select a route map, the default action is to allow all multicast traffic on the bridge domain; however, if you select a route map the default action changes to deny any traffic not explicitly matched to a filter entry in the route map.

Step 10 If your multicast source is in one site and is not stretched to the other sites, enable intersite multicast source option on the EPG.

Once you have enabled L3 Multicast on the BD, you must also enable multicast on the EPGs (part of multicast-enabled BDs) where multicast sources are connected.

- a) Select the EPG for which you want to enable Layer 3 multicast.
- b) In the right-hand sidebar, check the **Intersite Multicast Source** checkbox.



CHAPTER 22

QoS Preservation Across IPN

- [QoS and Global DSCP Policy, on page 261](#)
- [DSCP Policy Guidelines and Limitations, on page 261](#)
- [Configuring Global DSCP Policy, on page 262](#)
- [Set QoS Level for EPGs and Contracts, on page 263](#)

QoS and Global DSCP Policy

Cisco ACI Quality of Service (QoS) feature allows you to classify the network traffic in your fabric and then to prioritize and police the traffic flow to help avoid congestion in your network. When traffic is classified within the fabric, it is assigned a QoS Priority Level, which is then used throughout the fabric to provide the most desirable flow of packets through the network.

This release of Nexus Dashboard Orchestrator supports configuration of QoS level based on source EPG or a specific Contract. Additional options are available in each fabric directly. You can find detailed information on ACI QoS in [Cisco APIC and QoS](#).

When traffic is sent and received within the Cisco ACI fabric, the QoS Level is determined based on the CoS value of the VXLAN packet's outer header. In certain use cases, such as multi-pod or remote leaf topologies, the traffic must transit an intersite network, where devices that are not under Cisco APIC's management may modify the CoS values in the packets. In these cases you can preserve the ACI QoS Level between parts of the same fabric or different fabrics by creating a mapping between the Cisco ACI QoS level and the DSCP value within the packet.

DSCP Policy Guidelines and Limitations

When configuring the global DSCP translation policy, the following guidelines apply.



Note If you plan to use the global DSCP translation policy along with SD-WAN integration, skip this chapter and see the [SD-WAN Integration, on page 285](#) chapter instead for all information including the full list of guidelines and limitations.

- Global DSCP policy is supported for on-premises sites only.
- When defining the global DSCP policy, you must pick a unique value for each QoS Level.

- When assigning QoS level, you can choose to assign it to a specific Contract or an entire EPG.

If multiple QoS levels could apply for any given traffic, only one is applied using the following precedence:

- Contract QoS level: If QoS is enabled in the Contract, the QoS level specified in the contract is used.
- Source EPG QoS level: If QoS level is not specified for the Contract, the QoS level set for the source EPG is used.
- Default QoS level: If no QoS level is specified, the traffic is assigned Level 3 QoS class by default.

Configuring Global DSCP Policy

When traffic is sent and received within a Cisco ACI fabric, it is prioritized based on the ACI QoS Level, which is determined based on the CoS value of the VXLAN packet's outer header. When traffic exits the ACI fabric towards an intersite network, for example in multi-pod and remote leaf topologies, the QoS level is translated into a DSCP value which is included in the outer header of the VXLAN-encapsulated packet.

This section describes how to define the DSCP translation policy for traffic entering or exiting ACI fabric. This is required when traffic must transit through non-ACI networks, where devices that are not under Cisco APIC's management may modify the CoS values in the transiting packets.

Before you begin

- You should be familiar with Quality of Service (QoS) functionality within ACI fabrics.
- QoS is described in more detail in [Cisco APIC and QoS](#).

Procedure

Step 1 Log in to your Nexus Dashboard and open the Nexus Dashboard Orchestrator service.

Step 2 Create a new Tenant Policy.

- From the left navigation pane, choose **Application Management > Tenant Policies**.
- On the **Tenant Policy Templates** page, click **Add Tenant Policy Template**.
- In the **Tenant Policies** page's right properties sidebar, provide the **Name** for the template.
- From the **Select a Tenant** dropdown, choose the tenant with which you want to associate this template.

All the policies you create in this template as described in the following steps will be associated with the selected tenant and deployed to it when you push the template to a specific site.

By default, the new template is empty, so you need to add one or more tenant policies as described in the following steps. Note that you don't have to create every policy available in the template – you can define one or more policies of each type to deploy along with this template. If you don't want to create a specific policy, simply skip the step that describes it.

Step 3 Create a QoS DSCP policy.

- From the **+Create Object** dropdown, select **QoS DSCP**.
- In the right properties sidebar, provide the **Name** for the policy.

- c) (Optional) Click **Add Description** and provide a description for the policy.
- d) Provide policy details.

- **Admin State** – Enables or disables the policy.
- **Advanced Settings** – Click the arrow next to this section to expand.

Choose the DSCP value for each ACI QoS level. Each dropdown contains the default list of available DSCP values. You must choose a unique DSCP value for each level.

- e) Repeat this step to create any additional QoS DSCP policies.

Typically, we recommend applying this policy consistently across all sites that are part of your Multi-Site domain.

Step 4 Assign the policy to one or more sites.

- a) In the Fabric Policies template view, select **Actions > Sites Association**.
- b) In the **Associate Sites to <template>** dialog, select one or more sites for this policy template and click **Ok**.
- c) In the Fabric Policies template view, click **Deploy**.

After you save and deploy, the DSCP policy settings will be pushed to each site. You can verify the configuration by logging in to the site's APIC and navigating to **Tenants > infra > Policies > Protocol > DSCP class-CoS translation policy for L3 traffic**.

What to do next

After you have defined the global DSCP policy, you can assign the ACI QoS Levels to EPGs or Contracts as described in [Set QoS Level for EPGs and Contracts, on page 263](#).

Set QoS Level for EPGs and Contracts

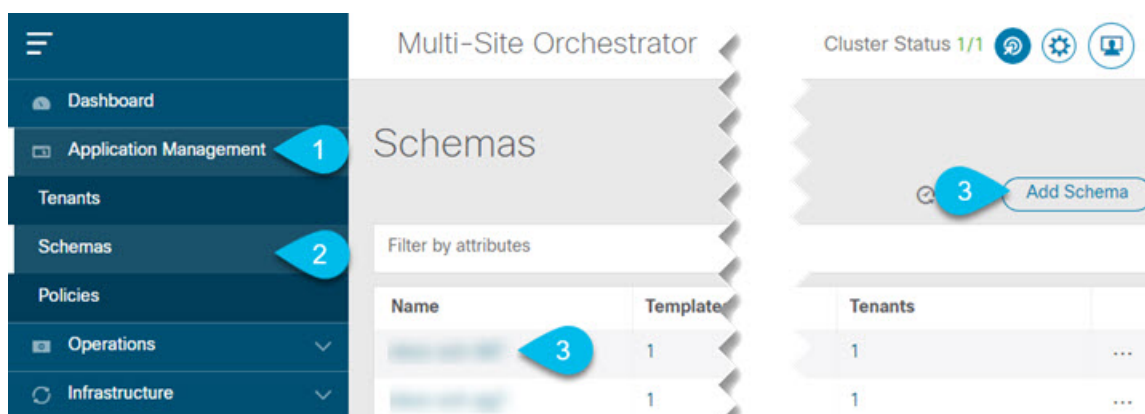
This section describes how to choose an ACI QoS level for traffic in your fabrics. You can choose to specify QoS for individual Contracts or entire EPGs.

Before you begin

- You must have defined the global DSCP policy, as described in [Configuring Global DSCP Policy, on page 262](#).
- You should be familiar with Quality of Service (QoS) functionality within ACI fabrics.
QoS is described in more detail in [Cisco APIC and QoS](#).

Procedure

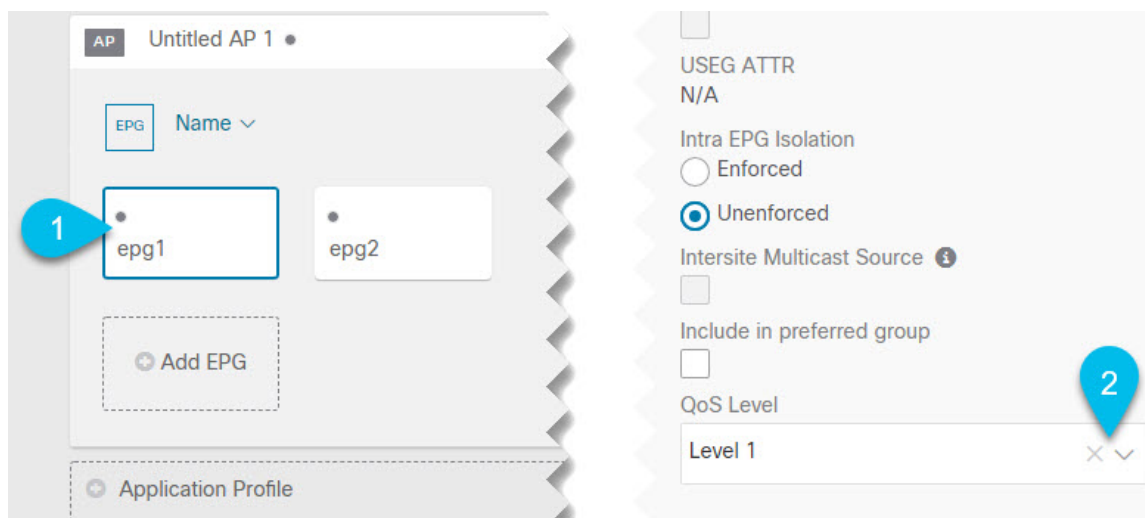
-
- Step 1** Log in to your Cisco Nexus Dashboard Orchestrator GUI.
 - Step 2** Choose the Schema you want to edit.



- a) Navigate to **Application Management > Schemas > .**
- b) Click the name of the schema you want to edit or **Add Schema** to create a new one.

The **Edit Policy** window will open.

Step 3 Pick a QoS Level for an EPG



- a) In the main pane, scroll down to the **EPG** area and select an EPG or click **Add EPG** to create a new one.
- b) In the right sidebar, scroll down to the **QoS Level** dropdown and choose the QoS Level you want to assign to the EPG.

Step 4 Pick a QoS Level for an EPG

The screenshot displays the configuration interface for Contracts and EPGs. The main pane on the left shows the 'CONTRACT' section with a search bar containing 'c1' and a '+ Filter' button. Below it is the 'VRF' section with a search bar containing 'vrf1' and a '+ Filter' button. The right sidebar contains the 'Filter Chain' section with a table showing a filter named 'f1' with a 'none' directive. Below this is the 'Service Graph' section with a dropdown menu. The 'ON-PREMISES PROPERTIES' section is highlighted, and the 'QoS Level' dropdown is set to 'Level 1'.

Name	Directive
f1	none

Service Graph

Select or find an item here

ON-PREMISES PROPERTIES

QoS Level

Level 1

- In the main pane, scroll down to the **Contract** area and select a Contract or click the + icon to create a new one.
- In the right sidebar, scroll down to the **QoS Level** dropdown and choose the QoS Level you want to assign to the Contract.



CHAPTER 23

SD-Access and ACI Integration

- [Cisco SD-Access and Cisco ACI Integration](#), on page 267
- [Macro Segmentation](#), on page 268
- [Cisco SD-Access and Cisco ACI Integration Guidelines](#), on page 270
- [Onboarding the DNA Center](#), on page 271
- [Configuring Connectivity Toward the SD-Access Domain](#), on page 272
- [Viewing the Status of the SD-Access to ACI Integration](#), on page 274
- [Extending a Virtual Network](#), on page 276
- [Mapping or Unmapping a VN to a VRF](#), on page 277
- [Configuring Transit Routing](#), on page 279

Cisco SD-Access and Cisco ACI Integration



Note Cisco Nexus Dashboard and Cisco DNAC integration allows for automation of a subset of network connectivity and macro segmentation scenarios across Nexus and campus SD-Access fabric deployments. This integration is under limited availability. Please contact your Cisco representative for additional information.

Cisco Software-Defined Access (SD-Access or SDA) is a solution within the Cisco Digital Network Architecture (DNA), which defines a campus-and-branch architecture that implements Cisco's Intent-Based Networking (IBN) framework. Cisco SD-Access defines a uniform policy-based wired and wireless network fabric that meets business needs with security, automation, and assurance. The Cisco Digital Network Architecture Controller (DNAC), in combination with Cisco Identity Services Engine (ISE), is the unified point of automation and management for the Cisco SD-Access fabric.

Release 3.7(1) of Cisco Nexus Dashboard Orchestrator (NDO) adds support for Cisco SD-Access and Cisco ACI integration. The purpose of SD-Access and ACI integration is to securely connect the campus-and-branch network to the data center network. With Release 3.7(1), NDO can perform the following functions:

- gather network and resource information from both domains
- automatically configure the VRF-Lite inter-domain connection at the ACI side
- provide the configuration of the next-hop device connected to the SD-Access border nodes
- provide cross-domain visibility

Macro Segmentation

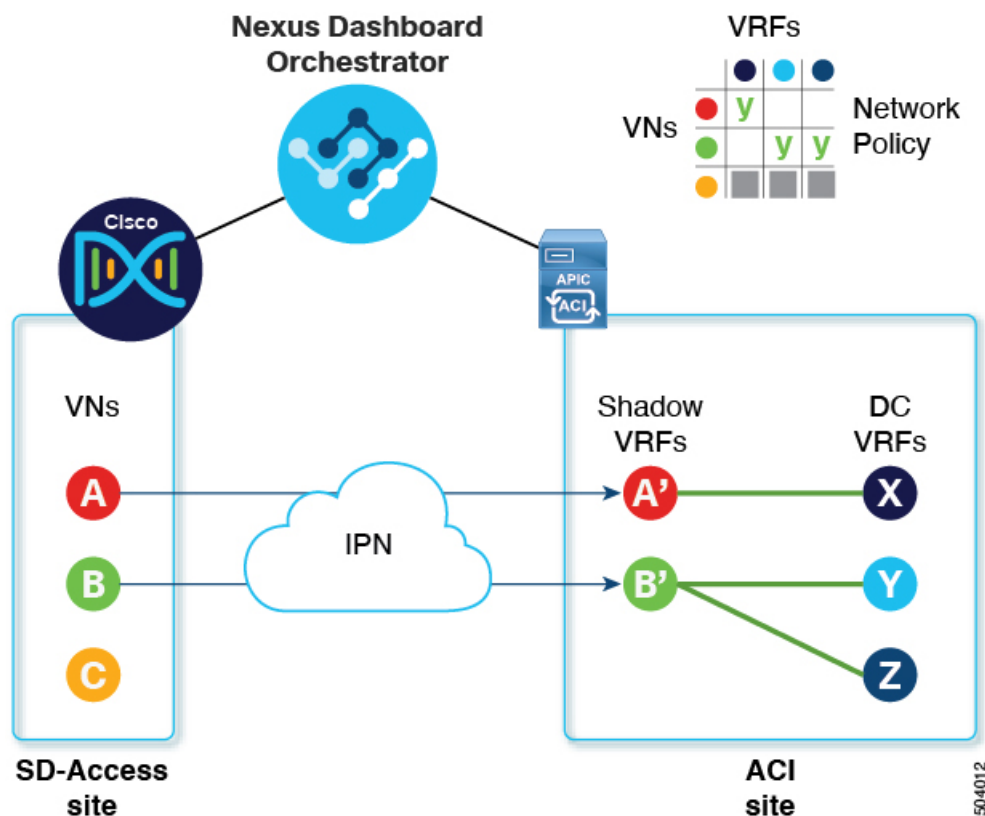
The Cisco SD-Access and Cisco ACI integration feature of Cisco Nexus Dashboard Orchestrator (NDO) allows macro segmentation of network elements between the ACI domain and the SD-Access domain.

In the ACI domain, entities such as EPGs, subnets, and VLANs are grouped as part of Virtual Routing and Forwarding instances (VRFs). When a VRF requires external communication, the VRF is associated with an IP interface (L3Out) on an ACI Border Leaf (BL). In the SD-Access domain, entities such as users, subnets, and IP pools can be grouped as Virtual Networks (VNs). When a VN requires external communication, the VN is associated with an SD-Access Border Node (BN) interface for IP handoff. The border interfaces of the two domains, ACI and SD-Access, can be physically connected through an IP network (IPN), but this basic connection does not provide connectivity between VRFs and VNs. With Cisco Nexus Dashboard Orchestrator Cisco SD-Access and Cisco ACI integration, an administrator can create policies to map (or "stitch") VRFs to VNs.

Macro Segmentation Workflow

A typical Cisco SD-Access and Cisco ACI integration workflow consists of the following steps, which refer to the figure below.

Figure 23: Macro Segmentation using NDO for SD-Access-to-ACI integration



- In an existing SD-Access site, a Cisco Digital Network Architecture Controller (DNAC) administrator has configured a campus fabric in which some entities require external access, such as access to the data center. The DNAC administrator has performed the following tasks:

- created virtual networks (VNs)
- associated IP address pools to those VNs
- configured L3 border nodes and associated interfaces
- created an IP (layer 3) handoff transit network
- configured Layer-3 handoff for those VNs that need external connectivity

Note that these tasks are normal DNAC administrative tasks and no special modification has been made for Cisco SD-Access and Cisco ACI integration.

- An NDO operator logs into and onboards the DNAC, using the DNAC credentials.

In the onboarding process, NDO automatically accesses the DNAC's REST API to query for sites, VNs, and border node devices. Upon discovering these entities, NDO learns which VNs are configured for external connectivity (L3 handoff) and on which Cisco SD-Access border nodes, and it learns their subnets. In the example shown in [Figure 23: Macro Segmentation using NDO for SD-Access-to-ACI integration, on page 268](#), the VNs A and B have been configured for L3 handoff, and these VNs are available to be extended into the ACI site. VN C is not configured for L3 handoff, and will not be available to the ACI site.

NDO continues to periodically query the DNAC for any ongoing configuration changes in the SD-Access fabric.

- The NDO operator configures connectivity between one or more ACI sites to one or more SD-Access sites. This includes specifying ACI site border leaf switches and interfaces, and VLAN and IP pools to be used for VRF-Lite configuration on border leaf interfaces. For directly connected interfaces (no IPN), the VRF-Lite configuration is derived from the configuration provisioned by DNAC for the IP handoff on the SDA border nodes, and VLANs and IP addresses are not taken from these pools.

NDO generates and displays next-hop device configuration for extended SD-Access VNs. This configuration can be applied manually to IPN devices if needed. NDO does not provision the IPN devices.

- The NDO operator extends a VN into the datacenter, making the VN available for connection to VRFs within the ACI domain.

Extending a VN creates an internal representation (a mirrored "shadow VRF") of the VN to represent the VN on the ACI domain. In the example in Figure 1, shadow VRFs A' and B' are automatically created in the ACI site to represent extended SD-Access VNs A and B. These shadow VRFs are stretched across all sites and pods within the ACI domain that require connectivity with the SD-Access domain. NDO automatically creates a schema and template in which these shadow VRFs are configured. The auto-created schema and template appear in NDO, but are read-only. The template is associated to the 'common' tenant and is associated with all 'SDA-Connectivity' enabled sites.

- The NDO operator creates a network policy to map an extended SD-Access VN to a datacenter VRF or VRFs that the VN needs to access. This action is also called "VRF stitching." The datacenter VRFs can be part of different "App tenants", which implies that this integration by design allows you to establish inter-VRF connectivity (a functionality usually referred to as "shared services").

In the example in Figure 1, the network policy shown stitches extended SD-Access VN A (extended as VRF A') to datacenter VRF X, and VN B (extended as VRF B') to datacenter VRFs Y and Z.

As a result of this mapping, a security policy relationship allowing all traffic is automatically established between the external EPG of the L3Out associated to the extended SD-Access VN and the vzAny logical object representing the datacenter VRF. The application of this contract allows free connectivity between

all the subnets of the extended SD-Access VN and all the subnets of the datacenter VRF that have been explicitly configured to be leaked across VRFs.

Cisco SD-Access and Cisco ACI Integration Guidelines

- An ACI site and an SD-Access site can be connected indirectly, through an external IP network (IPN), or directly, with back-to-back connections from ACI border leafs to SD-Access border nodes.
 - If the sites are directly connected, the connectivity between the two domains is configured automatically, including both the control plane and data plane.
 - If the sites are connected using an IPN, the IPN devices must support VRF Lite. NDO and DNAC do not provision the IPN devices, but NDO provides a sample configuration that can be applied to the IPN devices directly connected to the ACI border leafs and to the SD-Access border nodes.
- When multiple sites exist in either domain, note the following guidelines:
 - An SD-Access site can use another SD-Access site (SDA transit) to connect to the ACI sites.
 - When multiple sites exist in the SD-Access (campus) domain, each campus site can connect directly to the datacenter domain (direct peering), or through an intermediate network that could be a generic IP network (IPN), or through another campus site (indirect peering).
 - In a Multi-Site deployment, each ACI fabric that requires direct or indirect connectivity with the SD-Access (campus) domain must deploy a local L3Out connection. If the ACI fabric is a Multi-Pod fabric, the L3Out connection can be deployed only in a pod or a subset of the pods that are part of the same fabric.
- M:N mapping of VNs to VRFs is supported, within the limits described in [Scalability of SD-Access and ACI Integration, on page 271](#).
- M:N mapping of SD-Access sites to ACI sites is supported, within the limits described in [Scalability of SD-Access and ACI Integration, on page 271](#).
- From the DNAC, NDO learns about all SD-Access (campus) VNs and their subnets. When a VN is extended into the ACI site, NDO assumes that all subnets of that extended VN are reachable from ACI border leafs. NDO periodically checks for the presence of these subnets on ACI border leafs. In the **Status** column of the **Integrations > DNAC > Virtual Networks** table for an extended VN, NDO reports the subnets that are not yet reachable.
- By default, when an extended VN is mapped to a DC VRF, the ACI site does not advertise transit routes to the VN. The NDO administrator controls which ACI subnets are leaked into the shadow VRF of the VN as follows:
 - BD subnets that are internal to the ACI VRF are leaked only if the subnets are configured with “Shared between VRFs”.



Note When an SD-Access VN is mapped to multiple ACI VRFs, only non-overlapping prefixes across all mapped ACI VRFs should be configured as “shared between VRFs”.

- External subnets learned from L3Outs configured in the ACI VRF are leaked only if the subnets are configured with “Shared Route Control” and if transit routing is enabled.

For detailed information, see [Configuring Transit Routing, on page 279](#).

- The SD-Access site cannot provide Internet connectivity to the ACI site.
- Automation of IPv6 connectivity is not supported.
- Multicast traffic is not supported between the domains.

Scalability of SD-Access and ACI Integration

- Only a single DNAC can be onboarded to your NDO for SD-Access and ACI integration.
- Multiple SD-Access (campus) sites are supported if managed by a single DNAC.
- Up to 2 ACI sites are supported for peering with SD-Access. Each ACI site can be a single Pod fabric or a Multi-Pod fabric.
- A virtual network (VN) can be mapped to a maximum of 10 ACI VRFs.
- Up to 32 virtual networks (VNs) from the SD-Access domain can be extended into the ACI domain.

Software Compatibility

The minimum software versions that support macro segmentation for SD-Access and ACI integration are listed in the following table.

Product	Supported Product Versions
NDO	3.7 and later releases
ACI	4.2 and later releases
DNAC	2.3.3 and later releases

Onboarding the DNA Center

This section describes how to configure a Nexus Dashboard Orchestrator (NDO) to log in to a DNA Center (DNAC). After logging in, NDO can import the SD-Access site configuration information necessary to create a network connection between the SD-Access domain and an ACI domain.

Procedure

-
- Step 1** Log in to your NDO.
- Step 2** From the left navigation pane, select **Integrations > DNAC**.
- Step 3** In the main pane, click **Add DNAC** to onboard a DNA Center.
The **Add DNAC** dialog box opens.
- Step 4** In the **Add DNAC** dialog box, perform the following steps:

- a) Enter a **Name** for the DNA Center.
- b) Enter the URL or IP address of the DNA Center as the **Device IP**.
- c) Enter a **Username** credential for logging in to the DNA Center.

Read-only access is sufficient.

- d) Enter a **Password** credential for logging in to the DNA Center.
- e) Enter the password again in **Confirm Password**.
- f) Click **Add**.

NDO automatically logs in to the DNAC through the REST API and queries for the configuration of virtual networks (VNs) and border node devices in the SD-Access domain controlled by the DNAC.

What to do next

- Configure connectivity from the ACI site to the SD-Access site or IPN.
- Create network policies to allow communication between VNs in the DNAC's SD-Access domain and VRFs in the ACI domain.

Configuring Connectivity Toward the SD-Access Domain

This section describes the infrastructure-level configuration performed on NDO for Cisco SD-Access to ACI integration. For each ACI fabric, you must select the border leaf nodes and their associated interfaces that will provide connectivity toward the Cisco SD-Access domain.

Before you begin

You must have onboarded the DNA Center.

Procedure

-
- Step 1** Log in to your Nexus Dashboard Orchestrator.
 - Step 2** From the left navigation pane, select **Integrations > DNAC**.
 - Step 3** In the main pane, click the **Overview** tab.
A dashboard of DNA Center appears.
 - Step 4** On the right side of the **DNAC Details** box, click the link for **Configuring Connectivity**.
The **Fabric Connectivity Infra** page appears.
 - Step 5** From the left navigation pane, under **Sites**, select the ACI site to be connected.
An **Site Connectivity** pane appears on the right.
 - Step 6** From the **Site Connectivity** pane, scroll down to the **SDA Connectivity** control and set it to **Enabled**.
Several fields appear below the **SDA Connectivity** control. Configure the settings in the following substeps.

- a) From the **External Routed Domain** drop-down list, choose the external routed domain (L3 domain) to be connected.
This routed domain must be already defined on APIC.
- b) In the **VLAN Pool** field, enter a range of VLAN numbers.

A VLAN number from this pool will be assigned to the sub-interfaces or SVIs when extending a campus VN to the datacenter. The VLAN pool must be the same as, or a subset of, the VLAN pool associated to the external routed domain you selected in the previous step.

If the ACI to SD-Access connection is back-to-back, with no IPN, the VLAN ID is not assigned from this pool. Instead, the VLAN ID is determined by what has been provisioned by DNAC for the IP handoff on the SD-Access border nodes.
- c) Under **VRF Lite IP Pool Ranges**, click the + symbol next to **Add VRF Lite IP Pool Range** and enter an IP subnet in the **IP Address** field.

IP addresses from this subnet will be assigned to the sub-interfaces or SVIs when extending a campus VN to the datacenter.

If the ACI to SD-Access connection is back-to-back, with no IPN, these pools are not used. In this case, the IP addresses for the sub-interfaces are determined by what has been provisioned by DNAC for the IP handoff on the SD-Access border nodes.

Step 7

In the center pane that displays the pods of the ACI site, click **Add Leaf Node** under the pod that will connect to the SD-Access site.

A **Select a Leaf** pane appears on the right. Configure the settings in the following substeps.

- a) From the **Leaf Node** drop-down list in the **Select a Leaf** pane, choose the border leaf switch that will connect to the SD-Access domain.
- b) In the **Router ID** field, enter the border leaf router ID.
- c) Under Interfaces, click the + symbol next to **Add Interface**

The **Add Interface** dialog box appears.
- d) Enter the **Interface ID**.
- e) From the **Interface Type** drop-down list, select either **Sub-Interface** or **SVI**.
- f) Enter the **Remote Autonomous System Number**.

If the ACI to SD-Access connection uses an IPN, this number should match the ASN of the IPN.

If the ACI to SD-Access connection is back-to-back, with no IPN, this number should match the ASN of the SD-Access border nodes.

- g) Click **Save**.

Step 8

In the top bar of the **Fabric Connectivity Infra** page, click **Deploy**.

At this point, the configuration is not yet pushed to APIC. When the first VN is extended, the SD-Access connectivity is configured automatically.

Viewing the Status of the SD-Access to ACI Integration

The **Integrations > DNAC** menu displays details about the integration status and provides an inventory of available virtual networks (VNs).

Overview Tab

The **Overview** tab displays the following information windows:

- **DNAC Details:** Displays the overall status, IP address, and version of the connected DNAC. This window also contains a link to **Configure Connectivity**.
- A summary graphics dashboard for the following resources:
 - **DNAC Enabled Sites:** The number and type of SD-Access sites under management by the DNAC. The supported site types are on-premises, AWS, and NDFC.
 - **Virtual Networks:** The number of available VNs, and how many are extended or not extended.
 - **DC VRFs:** The number of datacenter VRFs available for sharing, and whether they are mapped or unmapped.

Virtual Networks Tab

Click the **Virtual Networks** tab to display details about the VNs.

The top window of the page repeats the summary graphics information from the **Overview** tab.

The **Virtual Networks** window of the page lists the virtual networks (VNs) that have been configured by DNAC for IP handoff on the SD-Access border nodes. A table of VNs displays the following information for each VN:

- **Status:** The current integration status of the VN, along with a color-coded icon indicating the severity of the status. The states are listed in the following table.

Status	Icon color (severity)	Description
Discovered	Green (Normal)	VN is discovered on SDA Border Nodes.
InProgress	Grey (Informational)	Reading the latest status of the VN after a configuration change. This is a temporary state. Tip You can click the Refresh icon in the upper right corner of the page to force an immediate polling of the status.
Success	Green (Normal)	VN is successfully extended.
BGPSessionIssues	Yellow (Warning)	BGP sessions are not established on all interfaces. Check each DC border leaf status for details.
RouteLeakPartial	Yellow (Warning)	VN subnets are partially propagated to the DC border leaf nodes. Check each DC border leaf status for details.

Status	Icon color (severity)	Description
RouteLeakNone	Red (Failure)	VN subnets are not yet propagated to the DC border leaf nodes. Click DC Sites in the VN table to check DC border leaf interfaces for issues.
MapVRFConfigFailure	Red (Failure)	Configuration failed on mapped VRFs. Retry the mapping.
DCSiteConfigFailure	Red (Failure)	VN extension failed on DC sites. Unextend the VN and extend again.

Click the status icon of a VN to display a sidebar containing additional details that can be helpful in troubleshooting warnings and failures.

- **Name:** The name assigned to the VN by the DNAC administrator.
- **Extended:** Indicates whether the VN has been extended.
- **DC Mapped VRFs:** The number of datacenter VRFs to which the VN is mapped. Click this number to open a sidebar displaying the associated schema, template, and tenant of mapped datacenter VRFs.
- **DC Sites:** The number of datacenter sites to which the VN is mapped. Click this number to open a sidebar displaying details of the datacenter sites, including the border leaf interfaces, BGP peering status, and next-hop device information.



Tip For IPN-connected border leaf interfaces, in the sidebar under "Peer Device Configuration", click "Show Details" for a sample configuration of an IPN device connected to this site.

- **Campus Sites:** The number of campus sites associated to this VN. Click this number to open a sidebar displaying details of the campus sites, including the border node interfaces, BGP peering status, and next-hop device information.



Tip For IPN-connected border node interfaces, in the sidebar under "Peer Device Configuration", click "Show Details" for a sample configuration of an IPN device connected to this site.

- **... (actions icon):** Click the icon to access actions for this VN.

The available actions depend on the current status of the VN, but may include the following:

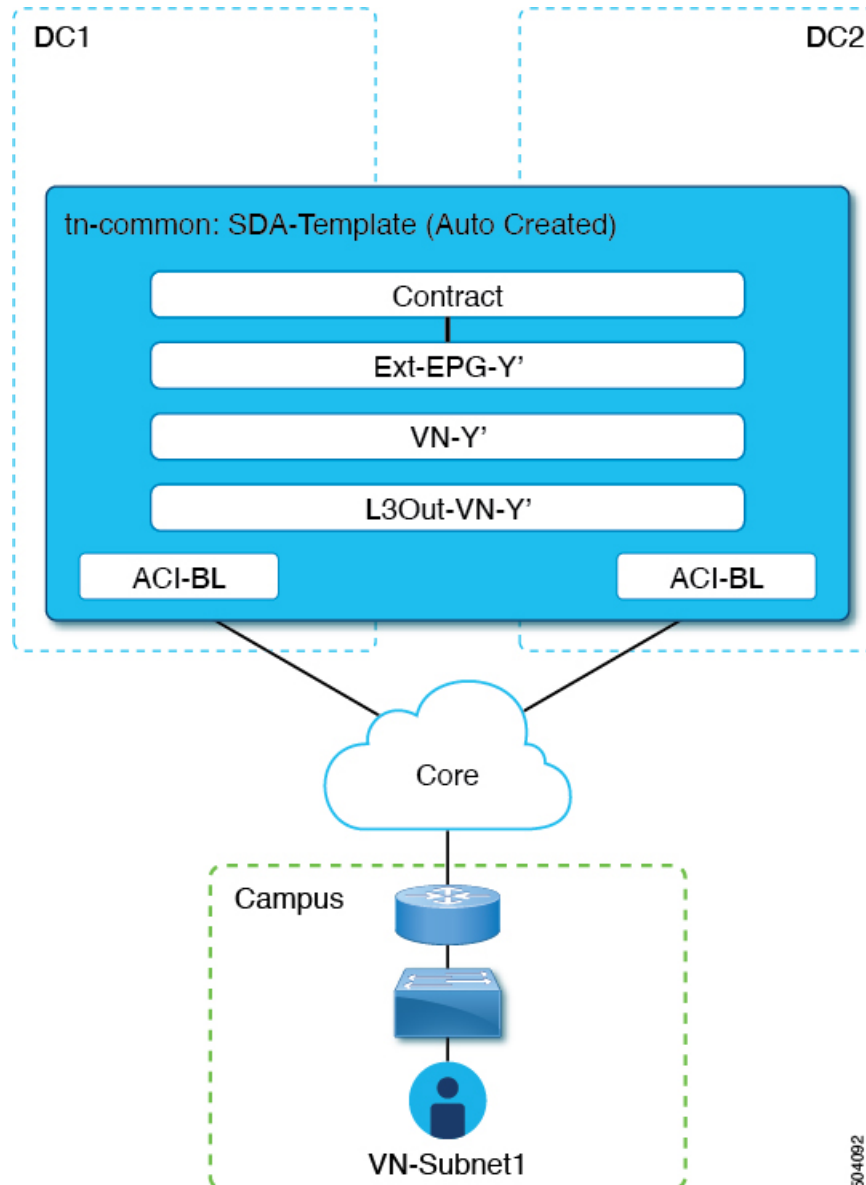
- Extend/Unextend the VN
- Map/Un-Map DC VRFs
- Enable/Disable Transit Route

The **Associated Templates** window of the **Virtual Networks** page appears when you have mapped a campus VN to a datacenter VRF.

Extending a Virtual Network

This section describes how to extend an SD-Access (campus) VN to the ACI (datacenter) fabrics. This action results in the creation of a VRF (and other associated configuration objects shown in [Figure 24: Extending a VN, on page 276](#)) representing the mirrored image of the campus VN on the DC side. The created objects are defined in an auto-generated template that is associated to the 'common' tenant.

Figure 24: Extending a VN



Before you begin

- You must have onboarded the DNA Center (DNAC).

- You must have configured connectivity to the SD-Access domain at the ACI site level.

Procedure

-
- Step 1** Log in to your Nexus Dashboard Orchestrator.
- Step 2** From the left navigation pane, select **Integrations > DNAC**.
- Step 3** In the main pane, click the **Virtual Networks** tab.
- A table of Virtual Networks (VNs) appears, displaying all VNs that have been configured by DNAC for IP handoff on the SD-Access border nodes.
- Step 4** In the row of the VN to be extended, click the actions menu (...) and select **Extend**.
- A dialog box opens, displaying the ACI sites and interfaces to which the VN will be extended. This information reflects the configuration settings in [Configuring Connectivity Toward the SD-Access Domain, on page 272](#).
- If you wish to revoke the extending of the VN later, click the actions menu (...) and select **Unextend**.
- Step 5** In the dialog box, click **Extend**.
- The VN is extended to all ACI sites where SD-Access connectivity is enabled, but it is not yet mapped to any ACI VRFs.
- Step 6**
-

What to do next

Verify the BGP Peering Status of the ACI border leaf interfaces:

- If the SD-Access border nodes and the ACI border leafs are connected directly (back-to-back), verify that, as a result of extending the campus VN, BGP sessions have been established between these devices. In **Integrations > DNAC > Virtual Networks**, click the **DC Sites** number to open a sidebar that displays details of the ACI border leaf interfaces. Check that the BGP Peering Status of the border leaf interfaces indicates "Up".
- If an IPN is deployed between the domains, retrieve the configuration samples to assist in configuring the next-hop devices that are directly connected to the SD-Access border nodes and to the ACI border leafs. In **Integrations > DNAC > Virtual Networks**, click the **DC Sites** number to open a sidebar that displays details of the ACI border leaf interfaces. For IPN-connected border leaf interfaces, click the "Show Details" link next to Peering Device Configuration to display a sample IPN device configuration. After configuring the IPN devices, check that the BGP Peering Status of the border leaf interfaces indicates "Up".

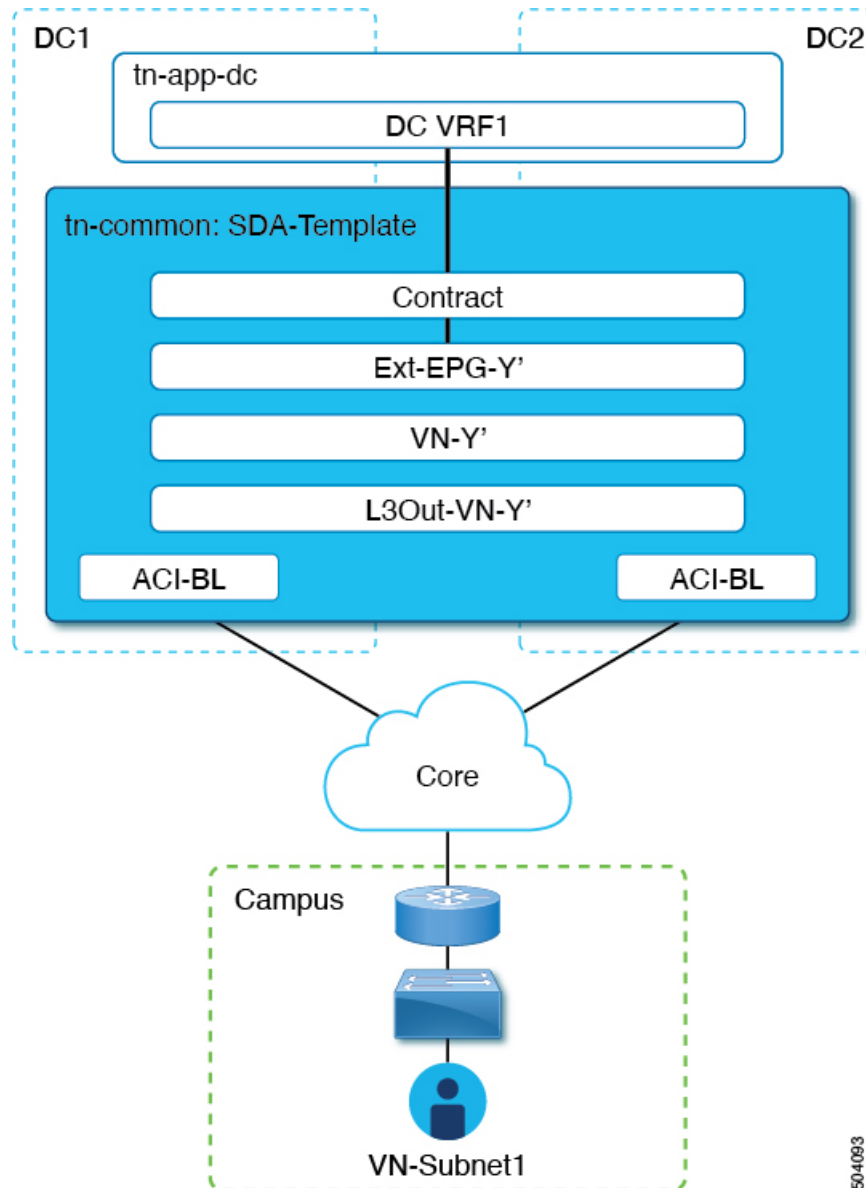
Map the extended VN to one or more ACI VRFs, as described in [Mapping or Unmapping a VN to a VRF, on page 277](#).

Mapping or Unmapping a VN to a VRF

This section describes how to map ("stitch") a virtual network (VN) to one or more datacenter (DC) VRFs in the ACI fabric. The mapping to a VRF results in the establishment of a contract relationship between the DC

VRF (represented by the "vzAny" object) and the external EPG previously provisioned in the 'common' tenant, as shown in [Figure 25: Mapping to a VRF](#), on page 278:

Figure 25: Mapping to a VRF



Before you begin

You must have extended the VN into the ACI site.

Procedure

Step 1 Log in to your Nexus Dashboard Orchestrator.

Step 2 From the left navigation pane, select **Integrations > DNAC**.

Step 3 In the main pane, click the **Virtual Networks** tab.

A table of Virtual Networks (VNs) appears, displaying all VNs that have been configured by DNAC for IP handoff on the SD-Access border nodes.

Step 4 In the row of the VN to be mapped, click the actions menu (...) and select **Map/Un-Map DC VRFs**.

A **Map/Un-Map DC VRFs** dialog box opens.

Step 5 In the **Map/Un-Map DC VRFs** dialog box, click the + icon next to **Add Mapped DC VRF**.

Step 6 From the drop-down list of VRFs, choose a VRF.

The selected VRF is added to a table that also displays the template for the VRF. Note the template name, as it will be needed in a later step.

If you wish to map the VN to additional VRFs, click the + icon again to choose additional VRFs from the drop-down list.

You can also un-map a DC VRF by deleting the existing mapping. To un-map a DC VRF, click the trash icon in the row of the VRF.

Step 7 Click **Save** and wait until the VN Status has changed to 'Success'.

Note

At this point, even if the VN Status indicates 'Success', data connectivity is not yet established between the extended VN and the DC VRF. The mapping operation has modified a template associated with the mapped VRF, and you must redeploy the template before connectivity is established. In the **Associated Templates** table under the VN table, the template associated with the mapped VRF appears.

Step 8 In the **Associated Templates** table in the **Integrations > DNAC > Virtual Networks** tab, click the link of the template associated with the mapped VRF.

The schema and template page opens.

Step 9 In the schema and template page, click **Deploy to sites**.

Step 10 If template review and approval (change control) is enabled, follow the change control workflow to redeploy the template. Otherwise, click **Deploy** to redeploy the template.

What to do next



Note

If you have unmapped a DC VRF, no template is displayed in the **Associated Templates** table. However, you must still go to **Application Management > Schemas** to re-deploy the associated template to remove the vzAny configuration. Otherwise, data plane communication remains enabled.

Configuring Transit Routing

When an extended SD-Access (campus) VN is mapped to an ACI (datacenter) VRF, any BD subnets of the DC VRF that are configured with “Advertised Externally” and “Shared between VRFs” flags are leaked into

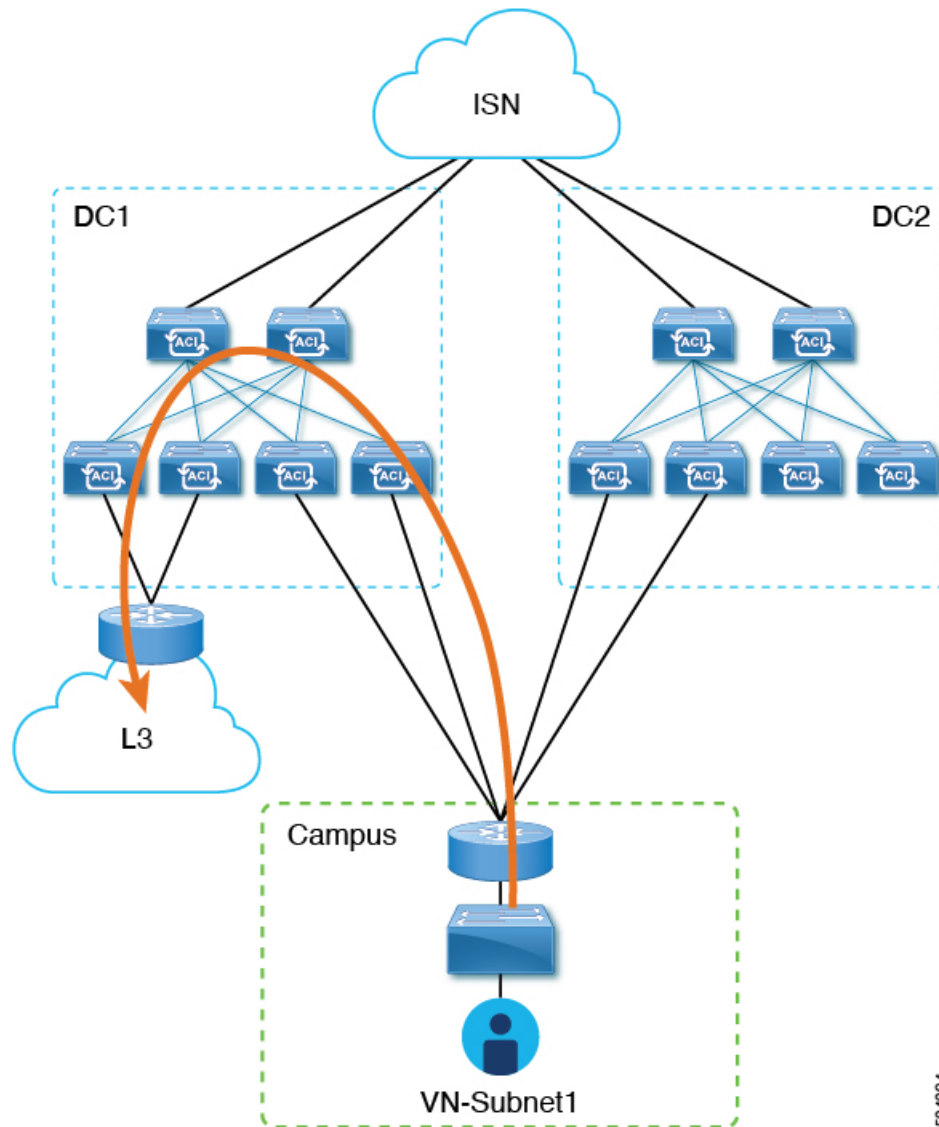
the ‘common’ tenant VRF and subsequently advertised toward the SD-Access domain. This ensures that campus users can gain access to the applications provisioned in the DC VRF.



Note When an SD-Access VN is mapped to multiple ACI VRFs, only non-overlapping prefixes across all mapped ACI VRFs should be configured as “shared between VRFs”.

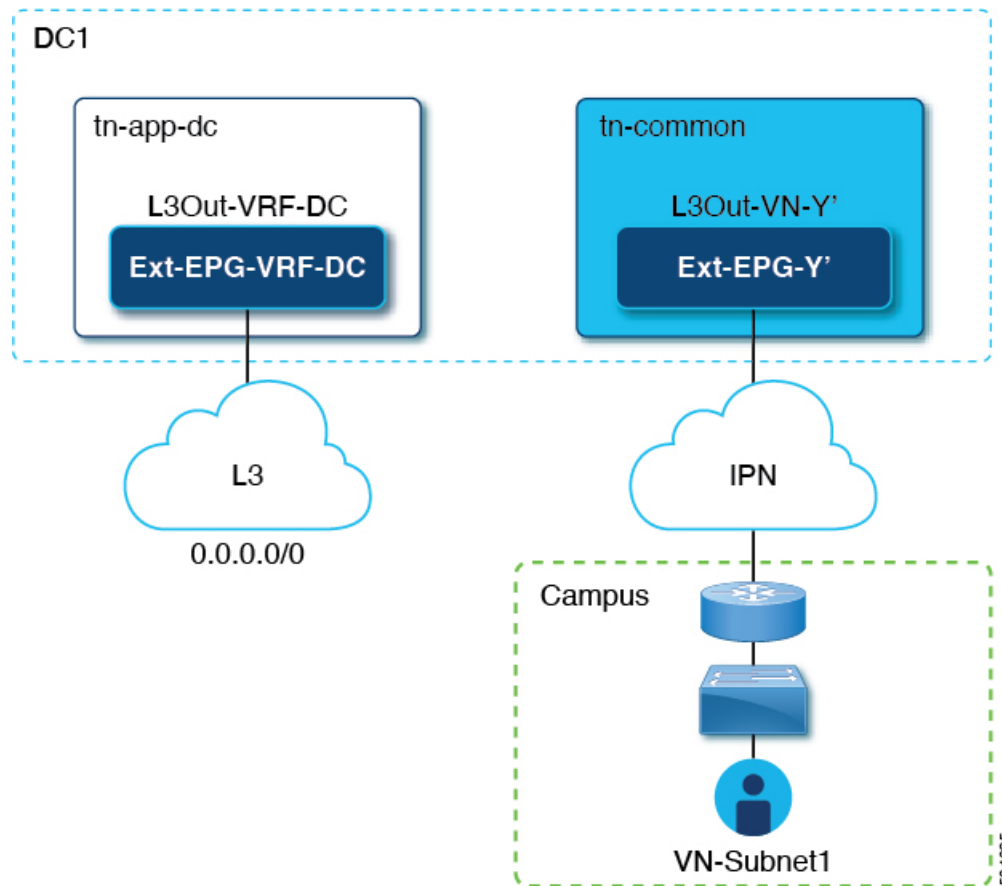
In addition to the advertisement of those BD subnets, there may be a requirement for campus users to access an external L3 network domain using the ACI domain as transit (Figure 26: ACI domain as transit, on page 280).

Figure 26: ACI domain as transit



In this scenario, an L3Out connection associated to the DC VRF (L3Out-DC-VRF) is usually provisioned for allowing connectivity to the external domain and external routes (a simple 0.0.0.0/0 default in the example in [Figure 27: L3Out connection, on page 281](#)) are imported in the DC VRF routing table (part of tn-app-dc).

Figure 27: L3Out connection



To ensure that campus users can connect to the external L3 domain through the datacenter, the external routes must be leaked to the tn-common VRF so that they can be advertised toward the campus domain through the L3Out connection (L3Out-VN-Y') auto-generated as a result of the campus VN extension to the DC.

Follow this procedure to enable leaking of the external routes:

Before you begin

You must have mapped an extended campus VN to a datacenter VRF and established connectivity.

Procedure

- Step 1** Log in to your Nexus Dashboard Orchestrator.
- Step 2** From the left navigation pane, select **Integrations > DNAC**.
- Step 3** In the main pane, click the **Virtual Networks** tab.
- Step 4** In the row of a successfully mapped campus VN, click the actions menu (...) and select **Enable Transit Route**.

This configuration ([Figure 28: Export Route Control, on page 282](#)) creates a 0.0.0.0/0 prefix under Ext-EPG-Y', with the following "Route Control" flags set that allow the advertising toward the IPN of all external routes leaked from the tn-app-dc tenant.

Figure 28: Export Route Control

Update Subnet 0.0.0.0/0

Subnet *
0.0.0.0/0

Route Control

☒ Export Route Control

☐ Import Route Control

☐ Shared Route Control

Aggregate

☒ Aggregate Export

External EPG Classification

☒ External Subnets for External EPG

☐ Shared Security Import

To disable transit routing, click the actions menu (...) and select **Disable Transit Route**.

Note

With either setting (enabled or disabled), the campus site has access to shared BD subnets internal to the ACI VRF.

Step 5 From the left navigation pane, choose **Application Management > Schemas** and navigate to the template for configuring the datacenter tenant application.

Step 6 In the datacenter tenant application template, configure the flags under the 0.0.0.0/0 prefix associated to Ext-EPG-VRF-DC of the DC VRF to be able to leak into tn-common the external routes learned from the Internet ([Figure 29: Shared Route Control, on page 282](#)).

Figure 29: Shared Route Control

Update Subnet 0.0.0.0/0

Subnet *
0.0.0.0/0

Route Control

☐ Export Route Control

☐ Import Route Control

☒ Shared Route Control

Aggregate

☒ Aggregate Shared Routes

External EPG Classification

☒ External Subnets for External EPG

☒ Shared Security Import

Note

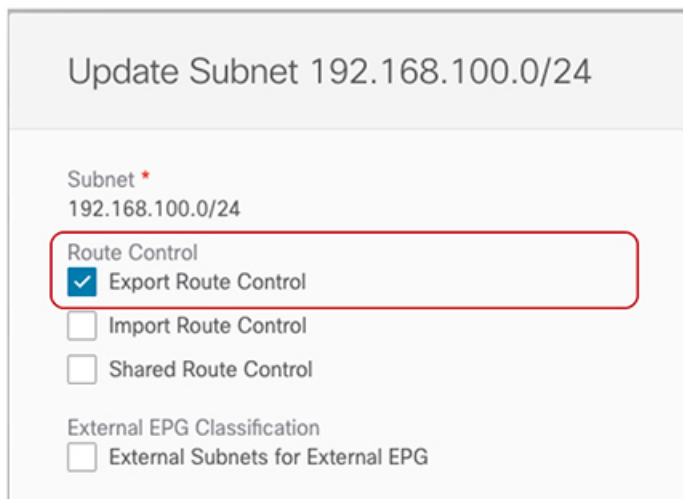
The setting shown ensures that all the external prefixes received on L3Out-VRF-DC are leaked to tn-common and are therefore advertised toward the campus domain. This setting also allows leaking of the 0.0.0.0/0 default route if it is received from the L3 domain. If desired, you can apply a more granular configuration where only a subset of the external

prefixes can be leaked to tn-common. This is achieved by creating specific entries matching those subsets of prefixes and applying to those entries the same flag configuration shown here.

Step 7 In the datacenter tenant application template, define a specific prefix under Ext-EPG-VRF-DC matching the campus VN subnet (or set of subnets) to be advertised toward the external L3 domain.

In the example shown in [Figure 30: Update Subnet, on page 283](#), this configuration is applied to a specific 192.168.100.0/24 prefix.

Figure 30: Update Subnet



Update Subnet 192.168.100.0/24

Subnet *
192.168.100.0/24

Route Control

- ☒ Export Route Control
- ☐ Import Route Control
- ☐ Shared Route Control

External EPG Classification

- ☐ External Subnets for External EPG

Note

Creating a separate prefix for a VN subnet provides the most granular level of control for the advertisement of campus VN subnets toward the external L3 domain. If such granular control is not needed, you can set the “Export Route Control” flags associated to the 0.0.0.0/0 prefix instead, which allows sending toward the external domain all the campus VN subnets that have been leaked into tn-app-dc from tn-common.



CHAPTER 24

SD-WAN Integration

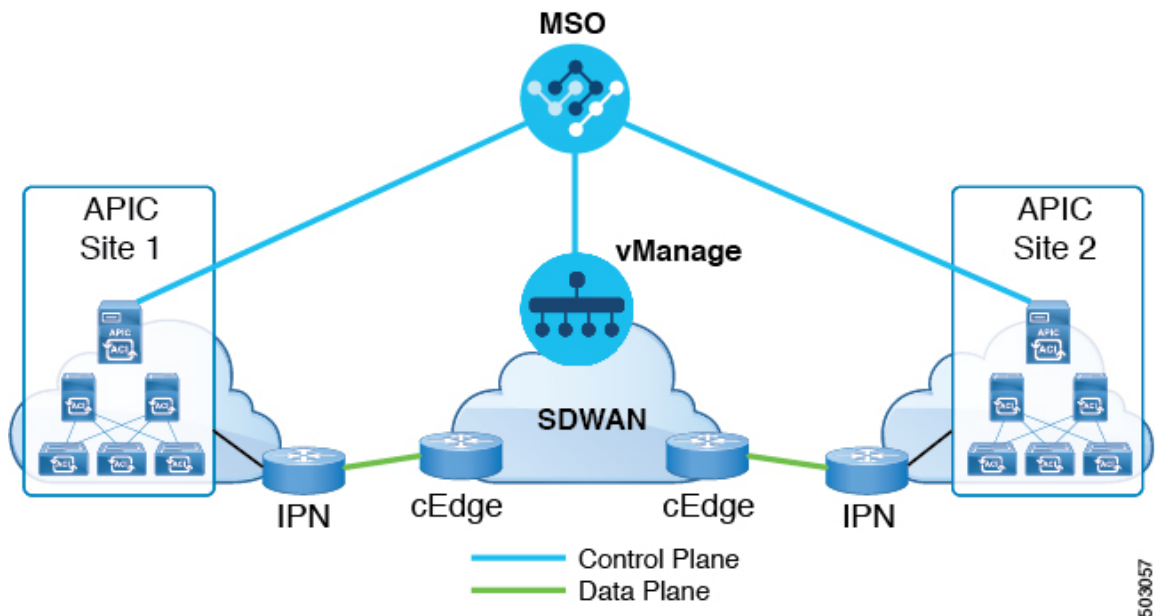
- [SD-WAN Integration, on page 285](#)
- [SD-WAN Integration Guidelines and Limitations, on page 286](#)
- [Adding a vManage Controller, on page 287](#)
- [Configuring Global DSCP Policy, on page 288](#)
- [Set QoS Level for EPGs and Contracts, on page 290](#)

SD-WAN Integration

Cisco Software-Defined Wide Area Network (SD-WAN) is a cloud-delivered overlay WAN architecture connecting branches to datacenter and multicloud environments through a single fabric. Cisco SD-WAN ensures predictable user experience for applications, optimizes SaaS, IaaS and PaaS connections, and offers integrated security either on-premises or in the cloud. Analytics capabilities deliver the visibility and insights necessary for you to isolate and resolve issues promptly and deliver intelligent data analysis for planning and what-if scenarios.

On the dataplane side, SD-WAN deploys an ASR or ISR routers as edge devices (shown as cEdge in the following diagram) with each fabric's spine switches connecting to these edge devices. SD-WAN is managed by a separate controller called vManage, which allows you to define service-level agreement (SLA) policies to determine how each packet's path within SD-WAN is chosen based on its DSCP value.

Figure 31: Multi-Site and SD-WAN Integration



Release 3.0(2) of Cisco Nexus Dashboard Orchestrator adds support for SD-WAN integration. You can configure the NDO to import SLA policies from a vManage controller, assign DSCP values to each SLA policy, and notify the vManage controller of the DSCP-to-SLA mapping. This enables you to apply preconfigured SLA policies to specify the levels of packet loss, jitter, and latency for intersite traffic over SD-WAN. The vManage controller, which is configured as an external device manager that provides SD-WAN capability, chooses the best possible WAN link that meets the loss, jitter, and latency parameters specified in the SLA policy.

Multi-Site SD-WAN integration allows traffic between multiple fabrics to traverse the SD-WAN network while enabling returning traffic from a remote site to retain the ACI QoS level assigned to it. After you register your Cisco NDO to vManage, it imports the SLA policies allowing you to translating the ACI QoS levels to the appropriate DSCP values. NDO then applies DSCP translation policy for traffic transiting SD-WAN to enable quality of service on the returning traffic.

Release 3.0(2) also enables you to assign ACI QoS levels to Contracts and EPGs directly in the NDO GUI. Any time traffic leaves the fabric, its QoS level is translated into a DSCP value, which vManage uses to pick a path for the traffic through SD-WAN.

SD-WAN Integration Guidelines and Limitations

When enabling Multi-Site and SD-WAN integration, the following guidelines apply.

- To enable uniform user QoS Level and DSCP translation for east-west traffic across sites with Multi-Site SD-WAN integration, the spine switches in each fabric must be connected to the SD-WAN edge devices, either directly or via multiple hops.

This is in contrast with the existing implementation of APIC SD-WAN integration for north-south traffic where the leaf switches must be connected to the SD-WAN edge devices.

- Global DSCP policy is supported for on-premises sites only.

- SD-WAN integration is supported for Nexus Dashboard Orchestrator deployments in Cisco Application Services Engine only.

For more information, see the [Deployment Overview](#) chapter in the *Cisco Nexus Dashboard Orchestrator Installation and Upgrade Guide*.

- When defining the global DSCP policy, you must pick a unique value for each QoS Level.
- In addition to existing DSCP policy values, you can import up to four SLA policies from vManage with one of the following values: 41, 42, 43, 45, 47 and 49.
- SLA policies must be already defined in your Cisco vManage.
- When assigning QoS level, you can choose to assign it to a specific Contract or an entire EPG.

If multiple QoS levels could apply for any given traffic, only one is applied using the following precedence:

- Contract QoS level: If QoS is enabled in the Contract, the QoS level specified in the contract is used.
- Source EPG QoS level: If QoS level is not specified for the Contract, the QoS level set for the source EPG is used.
- Default QoS level: If no QoS level is specified, the traffic is assigned Level 3 QoS class by default.

Adding a vManage Controller

This section describes how to add vManage controller to your Nexus Dashboard Orchestrator in order to import any configured SLA policies.

Procedure

Step 1 Log in to your Cisco Nexus Dashboard Orchestrator GUI.

Step 2 Add a vManage Controller.

- a) Navigate to **Integration > SD-WAN**.
- b) Click **Add Domain Controller** name.

The **Add Domain** window will open.

Step 3 Provide the vManage controller information.

In the **Add Domain** window that opens, provide the following details:

- Name of the vManage domain to display in your NDO.
- The device's fully qualified domain name or IP address.
- Username and password used to log in to the vManage controller.

Then click **Add** to save the vManage domain. After the vManage controller information is entered, it can take up to one min before the list of existing SLA policies is displayed in the main pane:

What to do next

Define the global DSCP policy in your Nexus Dashboard Orchestrator, as described in [Configuring Global DSCP Policy, on page 288](#)

Configuring Global DSCP Policy

When traffic is sent and received within a Cisco ACI fabric, it is prioritized based on the ACI QoS Level, which is determined based on the CoS value of the VXLAN packet's outer header. When traffic exits the ACI fabric from a spine switch towards an intersite network, the QoS level is translated into a DSCP value which is included in the outer header of the VXLAN-encapsulated packet.

This section describes how to define the DSCP translation policy for traffic entering or exiting ACI fabric. This is required when traffic must transit through non-ACI networks, such as between multiple fabrics separated by SD-WAN, where devices that are not under Cisco APIC's management may modify the CoS values in the transiting packets.

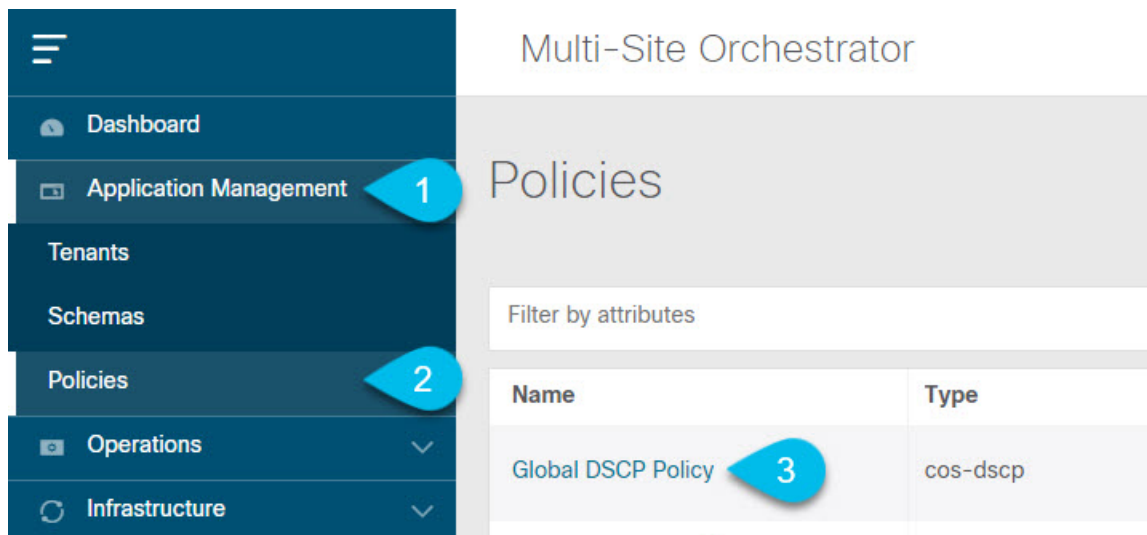
Before you begin

- You must have added a vManage controller to your NDO, as described in [Adding a vManage Controller, on page 287](#).
- You should be familiar with Quality of Service (QoS) functionality within ACI fabrics. QoS is described in more detail in [Cisco APIC and QoS](#).

Procedure

Step 1 Log in to your Cisco Nexus Dashboard Orchestrator GUI.

Step 2 Open the global DSCP policy configuration screen.



- Navigate to **Application Management > Policies**.
- Click **Global DSCP Policy** name.

The **Edit Policy** window will open.

Step 3 Update the global DSCP policy.

- a) Choose the DSCP value for each ACI QoS level.

Each dropdown contains the default list of available DSCP values as well as any values imported from the vManage SLA policies, for example `Voice-And-Video SLA (42)`.

- b) Choose the sites where you want to deploy the policy.

We recommend deploying the policy to all sites that are part of the Multi-Site domain in order to achieve a consistent end-to-end QoS behavior.

- c) Choose whether you want to enable the policy on each site when it is deployed.
d) Click **Save & Deploy**.

After you save and deploy, the DSCP policy settings will be pushed to each site. You can verify the configuration by logging in to the site's APIC and navigating to **Tenants > infra > Policies > Protocol > DSCP class-CoS translation policy for L3 traffic**.

What to do next

After you have defined the global DSCP policy, you can assign the ACI QoS Levels to EPGs or Contracts as described in [Set QoS Level for EPGs and Contracts, on page 290](#)

Set QoS Level for EPGs and Contracts

This section describes how to choose an ACI QoS level for traffic in your fabrics. You can choose to specify QoS for individual Contracts or entire EPGs.

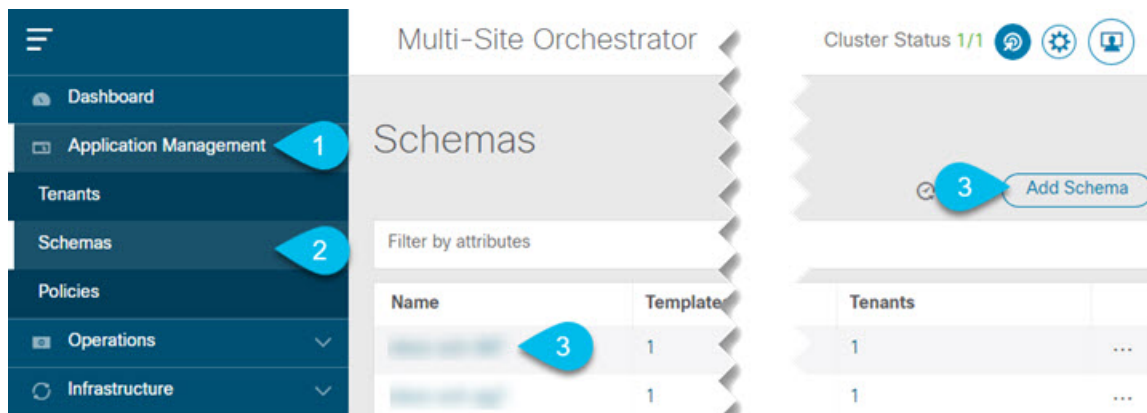
Before you begin

- You must have added a vManage controller to your NDO, as described in [Adding a vManage Controller, on page 287](#).
- You must have defined the global DSCP policy, as described in [Configuring Global DSCP Policy, on page 288](#).
- You should be familiar with Quality of Service (QoS) functionality within ACI fabrics.
QoS is described in more detail in [Cisco APIC and QoS](#).

Procedure

Step 1 Log in to your Cisco Nexus Dashboard Orchestrator GUI.

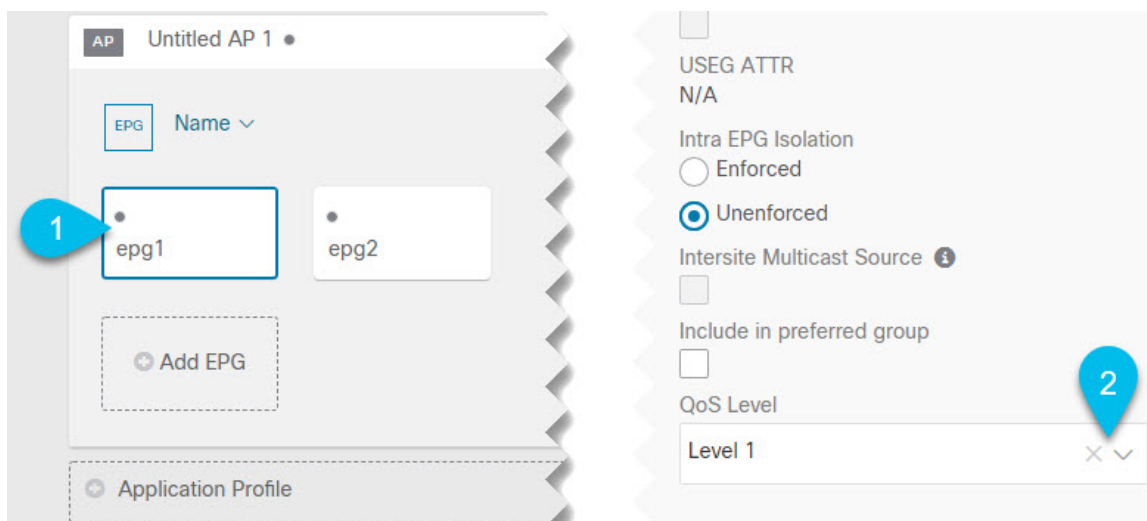
Step 2 Choose the Schema you want to edit.



- Navigate to **Application Management > Schemas > .**
- Click the name of the schema you want to edit or **Add Schema** to create a new one.

The **Edit Schema** window will open.

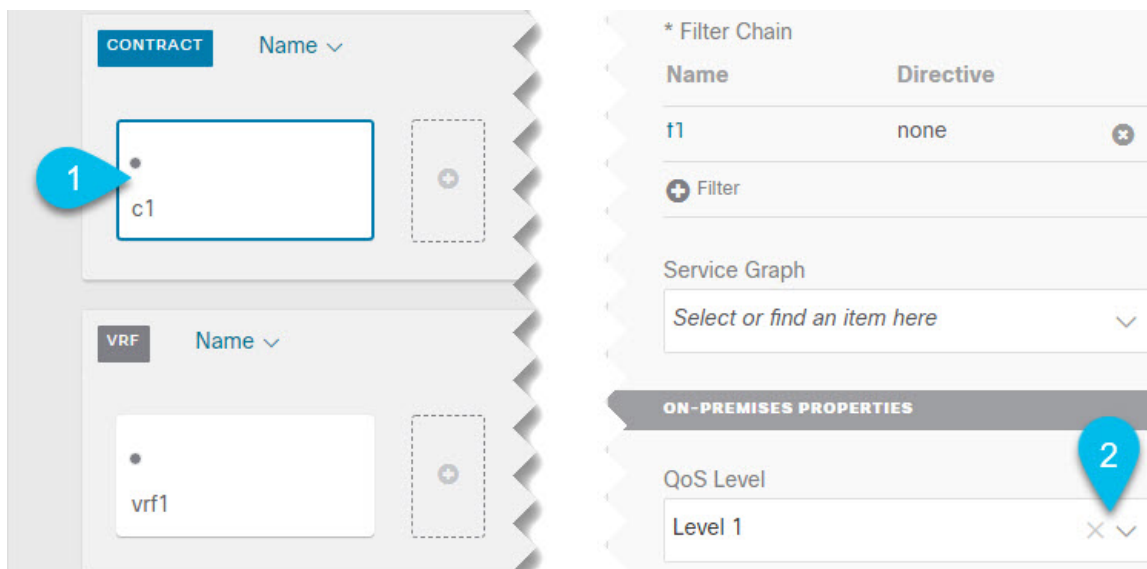
Step 3 Pick a QoS Level for an EPG



- In the main pane, scroll down to the **EPG** area and select an EPG or click **Add EPG** to create a new one.
- In the right sidebar, scroll down to the **QoS Level** dropdown and choose the QoS Level you want to assign to the EPG.

You must choose the QoS level based on the previously configured Global DSCP policy to ensure that intersite traffic from the EPG is treated with the desired SLA across the SD-WAN network.

Step 4 Pick a QoS Level for an EPG



- In the main pane, scroll down to the **Contract** area and select a Contract or click the + icon to create a new one.
- In the right sidebar, scroll down to the **QoS Level** dropdown and choose the QoS Level you want to assign to the Contract.

You must choose the QoS level based on the previously configured Global DSCP policy to ensure that intersite traffic between two EPGs is treated with the desired SLA across the SD-WAN network.



CHAPTER 25

Multi-Site and SR-MPLS L3Out Handoff

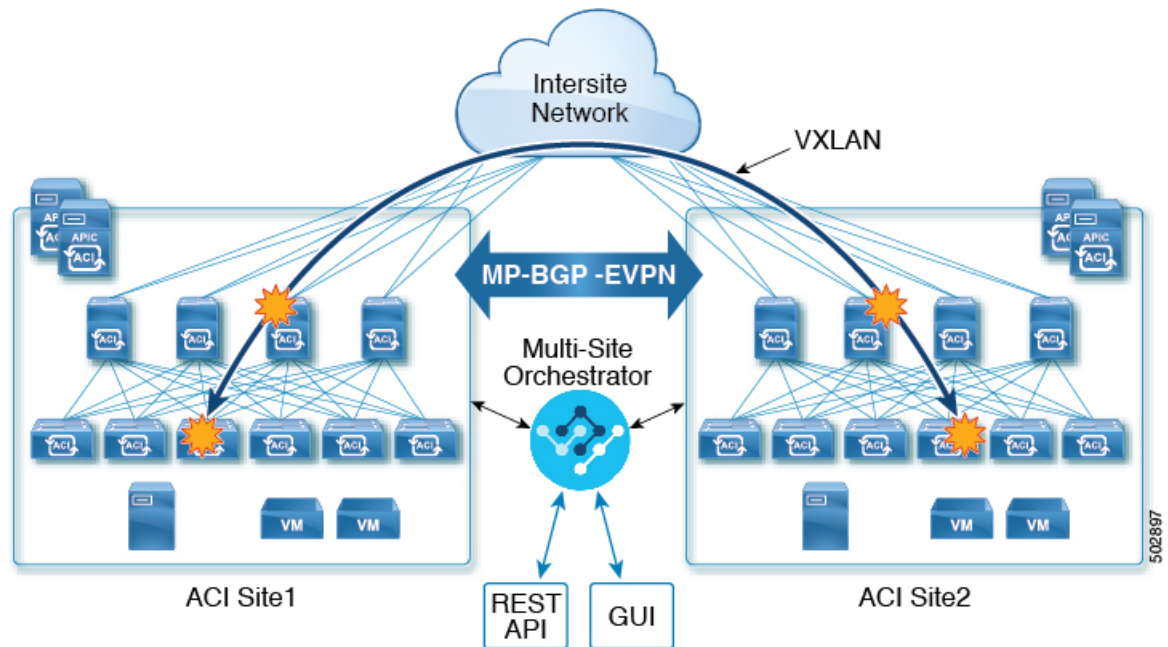
- [Overview and Use Cases, on page 293](#)
- [SR-MPLS Infra Requirements and Guidelines, on page 297](#)
- [SR-MPLS Tenant Requirements and Guidelines, on page 299](#)
- [Creating Custom QoS Policy for SR-MPLS, on page 301](#)
- [Creating SR-MPLS Infra L3Out, on page 303](#)
- [Creating SR-MPLS Route Map Policy, on page 306](#)
- [Configure EPG-to-External-EPG \(North-South\) Communication, on page 308](#)

Overview and Use Cases

Starting with Nexus Dashboard Orchestrator release 3.0(1) and APIC Release 5.0(1), the Multi-Site architecture provides better hand-off functionality between ACI border leaf (BL) switches and SR-MPLS networks.

In a typical Multi-Site deployment, traffic between sites is forwarded over an intersite network (ISN) via VXLAN encapsulation:

Figure 32: Multi-Site and ISN

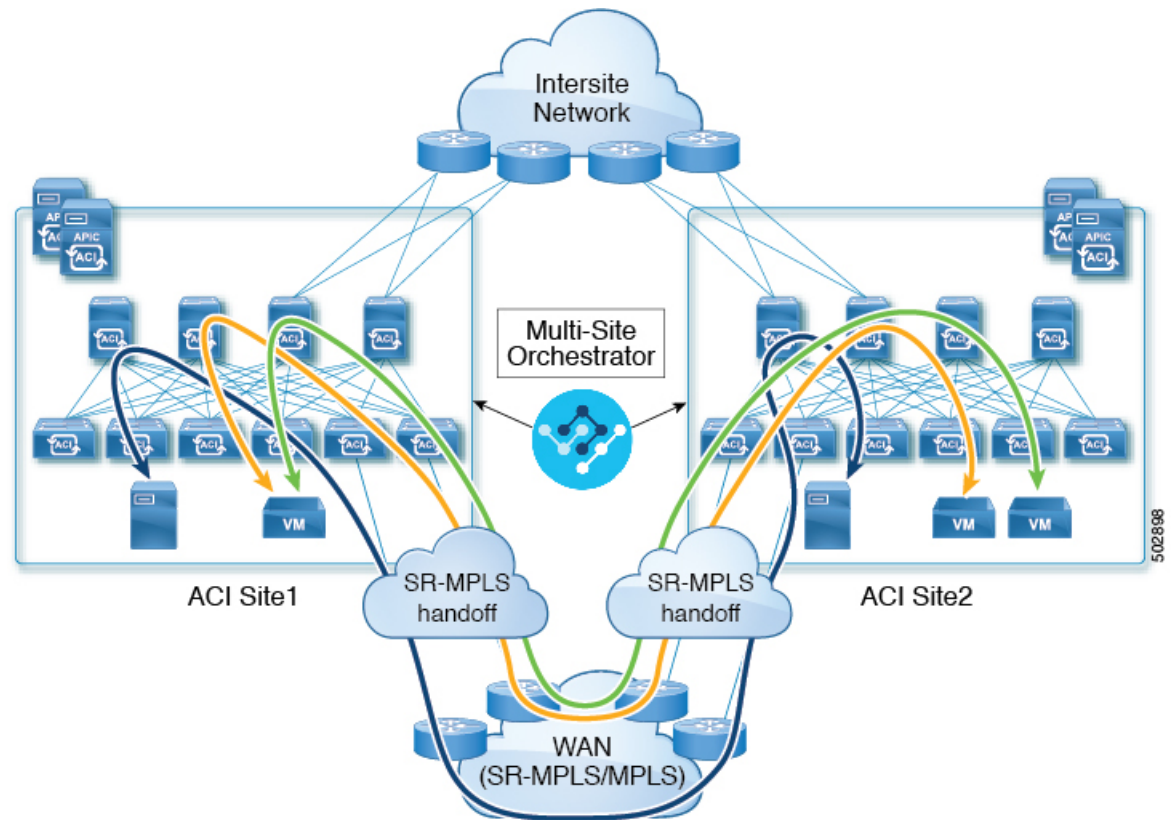


With Release 3.0(1), MPLS network can be used in addition to or instead of the ISN allowing inter-site communication via WAN, as shown in the following figure. In order to force East-West Layer 3 communication to follow the SR-MPLS L3Out data path (instead of the VXLAN data path across the ISN), several restrictions had to be applied to this SR-MPLS hand-off use case:

- The VRF to which the SR-MPLS L3Out belongs must not be stretched across sites.
- Because of the above restriction, every site must deploy one (or more) local SR-MPLS L3Outs for each defined site-local VRF.
- Contracts must not be applied between site-local EPGs belonging to different VRFs.

This forces the communication to follow the SR-MPLS L3Out data path.

Figure 33: Multi-Site and MPLS



Additional Use Cases in Release 4.0(2) and Later

Prior to release 4.0(2), if you wanted to deploy the SR-MPLS use case, you would define a special "SR-MPLS" template that could be associated with only a single site and not stretched across multiple sites. In this case, if you had two sites managed by your Nexus Dashboard Orchestrator and connected via an SR-MPLS network, and you wanted to establish communication between an EPG in `site1` and another EPG in `site2`, you had to deploy two separate SR-MPLS L3Outs (one in each site) associated with two separate VRFs and establish contracts between the EPG in each site and that site's SR-MPLS L3Out (instead of directly between the EPGs). In other words, the EPGs' traffic would always use the SR-MPLS L3Out data path even for EPG-to-EPG communication across sites without integrating with the traditional Multi-Site data plane for East-West traffic.

Beginning with release 4.0(2), the SR-MPLS L3Outs can function similar to the traditional IP-based L3Outs which allows you to use the SR-MPLS L3Out hand-offs exclusively for North-South connectivity between a site and an external network, while all East-West traffic is handled in the traditional Multi-Site manner using VXLAN-encapsulated data plane across the ISN. This means that the SR-MPLS hand-offs can now be treated as traditional IP-based hand-offs and the same VRF can deploy a mix of IP and SR-MPLS L3Outs. These changes add support for the following specific use cases:

- Centralized deployment in a site of an SR-MPLS L3Out which belongs to a specific VRF.

All traffic from endpoints that are part of that VRF (connected to the same site or to different sites) can leverage that centralized SR-MPLS L3Out for North-South connectivity. Note that this requires the VRF to be stretched across sites.

- Deployment of multiple sites each with their own local SR-MPLS L3Outs and intra-VRF traffic using the local L3Out if it is available or a remote SR-MPLS L3Out from another site (intersite L3Out).

In this case, the remote SR-MPLS L3Out can be used as a simple backup or to reach unique external prefixes received on the remote SR-MPLS L3Out. Traffic will transit from a local EPG to the local SR-MPLS L3Out and if that path is down or the route is unavailable it can take another site's remote SR-MPLS L3Out.

- Similar use cases are supported for shared services, where application EPG in one VRF can use SR-MPLS L3Out in a different VRF, either in the local or remote site.

In this case, the EPGs can be in a different tenant as well. For example, Tenant1 in Site1 can contain the application EPGs which will use an SR-MPLS L3Out in Tenant2 in Site2.

- Ability to combine IP-based and SR-MPLS hand-offs.
- SR-MPLS hand-off can now be used to connect to Provider Edge (PE) devices that are part of an SR-MPLS core network as well as PEs that are part of a regular MPLS network, in which case the hand-off can be considered as a simple MPLS hand-off.

Using SR-MPLS L3Outs (instead of traditional IP-based L3Outs) allows for operational simplification at higher scale by removing the need for the VRF-Lite configuration which requires creation of separate BL nodes, BL logical interfaces, and routing peering for each VRF that needs to be connected to the external network. With the SR-MPLS L3Outs, the logical nodes and logical interfaces are defined once in the infra tenant, together with a single MP-BGP EVPN peering with the external devices. This infra L3Out construct can then be used to provide external connectivity to multiple tenant VRFs and all the VRFs' prefixes are exchanged using the common MP-BGP EVPN control plane.

The following sections describe guidelines, limitations, and configurations specific to managing Schemas that are deployed to sites from the Nexus Dashboard Orchestrator. Detailed information about MPLS hand off, supported individual site topologies (such as remote leaf support), and policy model is available in the [Cisco APIC Layer 3 Networking Configuration Guide](#).

Configuration Workflow

Other sections in this document detail the required configurations, but in short you will go through the following workflow:

- Create an SR-MPLS QoS policy.

SR-MPLS Custom QoS policy defines the priority of the packets coming from an SR-MPLS network while they are inside the ACI fabric based on the incoming MPLS EXP values defined in the MPLS QoS ingress policy. It also marks the CoS and MPLS EXP values of the packets leaving the ACI fabric through an MPLS interface based on IPv4 DSCP values defined in MPLS QoS egress policy.

This is an optional step and if no custom ingress policy is defined, the default QoS Level (`Level3`) is assigned to packets inside the fabric. If no custom egress policy is defined, the default EXP value of 0 will be marked on packets leaving the fabric.

- Create an SR-MPLS Infra L3Out.

This configures an L3Out for traffic leaving a site that is connected to an SR-MPLS network.

- Create SR-MPLS route map policy.

Route maps are sets of `if-then` rules that enable you to specify which routes are advertised out of the Tenant SR-MPLS L3Out. Route maps also enable you to specify which routes received from the DC-PE routers will be injected into the BGP VPNv4 ACI control plane.

- If you want to deploy a use case similar to releases prior to release 4.0(2), create the VRF, SR-MPLS L3Out, and SR-External EPG for each site connected via an SR-MPLS network and establish a contract within each site between that site's tenant EPG and SR-External EPG.

In this case, all communication from one site will follow the North-South route egressing your Multi-Site domain towards the external SR-MPLS network. If the traffic is destined to an EPG in another site managed by your Orchestrator, it will ingress the other fabric from the external network using that site's SR-MPLS L3Out.

- If you want to use the SR-MPLS L3Outs in the same way as the standard IP-based L3Out exclusively for North-South communication, you can create the VRFs, SR-MPLS L3Outs, EPGs, and contracts as you typically would for all existing EPG-to-EPG communication use cases.

SR-MPLS Infra Requirements and Guidelines

If you want to use your Nexus Dashboard Orchestrator to manage SR-MPLS L3Out hand-offs for an ACI fabric connected to an SR-MPLS network:

- Any changes to the topology, such as node updates, are not reflected in the Orchestrator configuration until site configuration is refreshed, as described in [Refreshing Site Connectivity Information, on page 153](#).
- Tenants deployed to a site that is connected via an SR-MPLS network will have a set of unique configuration options specifically for SR-MPLS configuration. Tenant configuration is described in the "Tenants Management" chapter of the [Multi-Site Configuration Guide, Release 3.1\(x\)](#).
- Remote leaf switches are not supported as source or destination for any Multi-Site traffic flows.
- SR-External EPGs that are part of preferred group cannot be the providers of a shared service (inter-VRF) contract.
- Preferred Group is not supported for Intersite SR-MPLS L3Outs.
- vzAny is not supported as a shared service provider.
- VRF that is enabled for Preferred Group cannot be a vzAny consumer.
- We recommend configuring tenant contract objects under dedicated template to avoid circular dependencies with other configuration objects that use the same contracts
- We recommend keeping SR L3Outs and External EPGs in a dedicated template.

This will allow you to stretch other objects (such as VRFs, BDs, EPGs) to sites which do not have an infra SR L3Out.

- When using SR-MPLS L3Out instead of traditional IP-based L3Outs:
 - Host-based routing advertisement is not supported for bridge domains that are stretched across sites.
 - Tenant Routed Multicast (TRM) is not supported with SR-MPLS L3Outs, so they can only be used for establishing Layer 3 unicast communication with the external network domain.

Supported Hardware

The SR-MPLS hand-off is supported for the following platforms:

- **Border Leaf switches:** The "FX", "FX2", and "GX" switch models.
- **Spine switches:**
 - Modular spine switch models with "LC-EX", "LC-FX", and "GX" at the end of the linecard names.
 - The Cisco Nexus 9000 series N9K-C9332C and N9K-C9364C fixed spine switches.
- **DC-PE routers:**
 - Network Convergence System (NCS) 5500 Series
 - ASR 9000 Series
 - NCS 540 or 560 routers

SR-MPLS Infra L3Out

You will need to create an SR-MPLS Infra L3Out for the fabrics connected to SR-MPLS networks as described in the following sections. When creating an SR-MPLS Infra L3Out, the following restrictions apply:

- Each SR-MPLS Infra L3Out must have a unique name.

The SR-MPLS Infra L3Out allows you to establish the control plane and data plane connectivity between the ACI border leaf switches and the external Provider Edge (PE) devices. SR-MPLS L3Outs that belong to various tenant VRFs can then leverage that Infra L3Out connectivity to establish communication with the external network domain.
- You can have multiple SR-MPLS infra L3Outs connecting to different routing domains, where the same border leaf switch can be in more than one L3Out, and you can have different import and export routing policies for the VRFs toward each routing domain.
- Even though a border leaf switch can be in multiple SR-MPLS infra L3Outs, a border leaf switch/provider edge router combination can only be in one SR-MPLS infra L3Out as there can be only one routing policy for a user VRF/border leaf switch/DC-PE combination.
- If there is a requirement to have SR-MPLS connectivity from multiple pods and remote locations, ensure that you have a different SR-MPLS infra L3Out in each of those pods and remote leaf locations with SR-MPLS connectivity.
- If you have a multi-pod or remote leaf topology where one of the pods is not connected directly to the SR-MPLS network, that pod's traffic destined for the SR-MPLS network will use standard IPN path to another pod, which has an SR-MPLS L3Out. Then the traffic will use the other pod's SR-MPLS L3Out to reach its destination across SR-MPLS network.
- Routes from multiple VRFs can be advertised from one SR-MPLS Infra L3Out to provider edge (PE) routers connected to the nodes in this SR-MPLS Infra L3Out.

PE routers can be connected to the border leaf directly or through other provider (P) routers.
- The underlay configuration can be different or can be the same across multiple SR-MPLS Infra L3Outs for one location.

For example, assume the same border leaf switch connects to PE-1 in domain 1 and PE-2 in domain 2, with the underlay connected to another provider router for both. In this case, two SR-MPLS Infra L3Outs will be created: one for PE-1 and one for PE-2. But for the underlay, it's the same BGP peer to the provider router. Import/export route-maps will be set for EVPN session to PE-1 and PE-2 based on the corresponding route profile configuration in the user VRF.

MPLS Custom QoS Policies

Following is the default MPLS QoS behavior:

- All incoming MPLS traffic on the border leaf switch is classified into QoS Level 3 (the default QoS level).
- The border leaf switch will retain the original DSCP values for traffic coming from SR-MPLS without any remarking.
- The border leaf switch will forward packets with the default MPLS EXP (0) to the SR-MPLS network.

Following are the guidelines and limitations for configuring MPLS Custom QoS policies:

- Data Plane Policers (DPP) are not supported at the SR-MPLS L3Out.
- Layer 2 DPP works in the ingress direction on the MPLS interface.
- Layer 2 DPP works in the egress direction on the MPLS interface in the absence of an egress custom MPLS QoS policy.
- VRF level policing is not supported.

SR-MPLS Tenant Requirements and Guidelines

While the Infra MPLS configuration and requirements are described in the Day-0 operations chapter, the following restrictions apply for any user Tenants you will deploy to sites that are connected to SR-MPLS networks.

- In case when traffic between two EPGs in the fabric needs to go through the SR-MPLS network:
 - Contracts must be assigned between each EPG and the SR-External EPG defined on the local Tenant SR-MPLS L3Out.
 - If both EPGs are part of the same ACI fabric but separated by an SR-MPLS network (for example, in Multi-Pod or remote leaf cases), the EPGs must belong to different VRFs and not have a contract between them nor route-leaking configured.
 - If EPGs are in different sites, they can be in the same VRF, but there must **not** be a contract configured directly between the EPGs and any other remote EPG that is part of the same VRF.
- When configuring a route map policy for the SR-MPLS L3Out:
 - Each L3Out must have a single export route map. Optionally, it can also have a single import route map.
 - Routing maps associated with any SR-MPLS L3Out must explicitly define all the routes, including bridge domain subnets, which must be advertised out of the SR-MPLS L3Out.

- If you configure a `0.0.0.0/0` prefix and choose to not aggregate the routes, it will allow the default route only.
- However, if you choose to aggregate routes for the `0.0.0.0/0` prefix, it will allow all routes.
- You can associate any routing policy with any tenant L3Out.
- Beginning with Nexus Dashboard release 4.0(1), transit routing between SR-MPLS networks is supported using the same or different VRFs for fabrics running Cisco APIC release 5.1(1) or later.

Figure 34: Transit Routing Configuration Using Single VRF

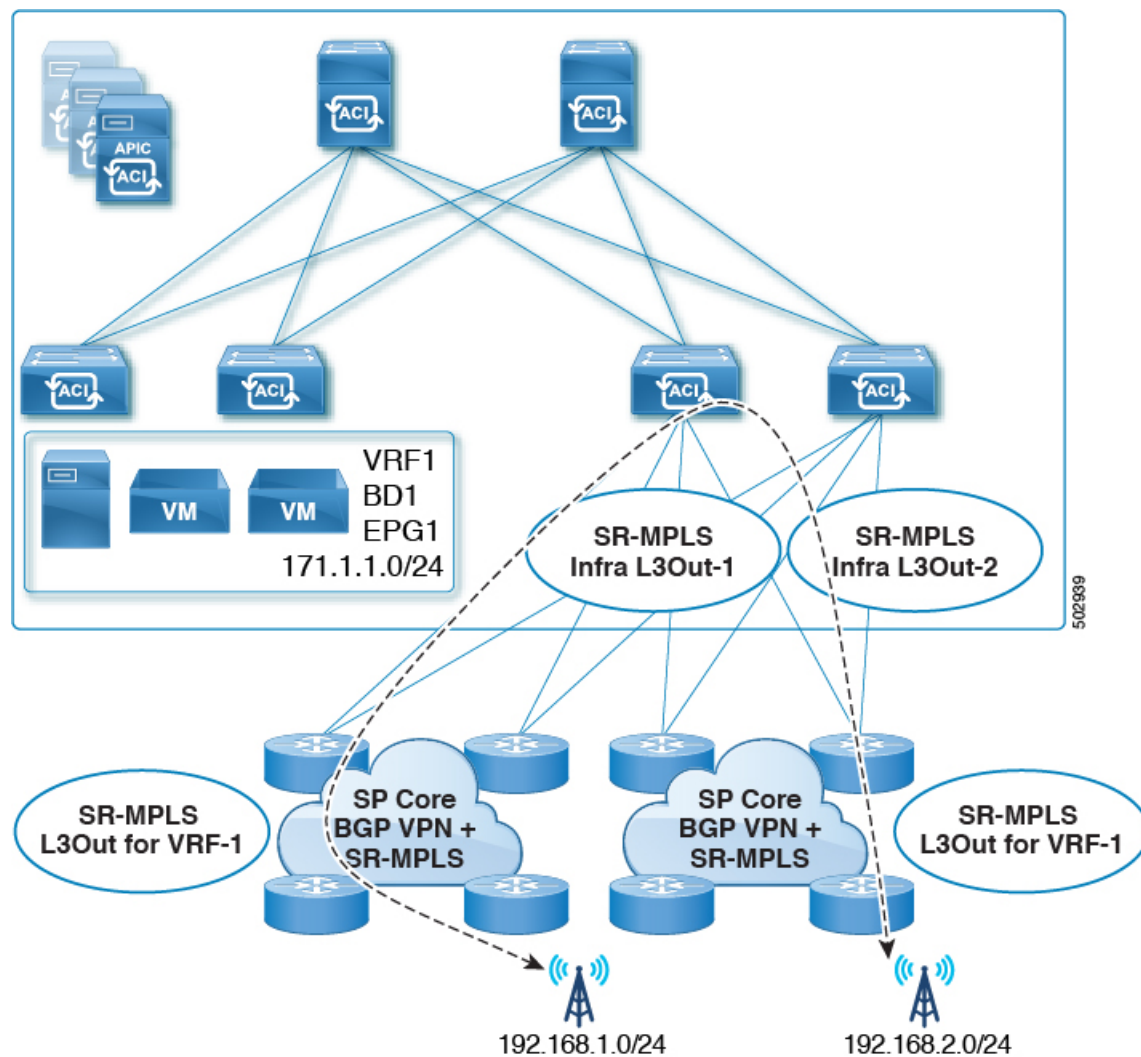
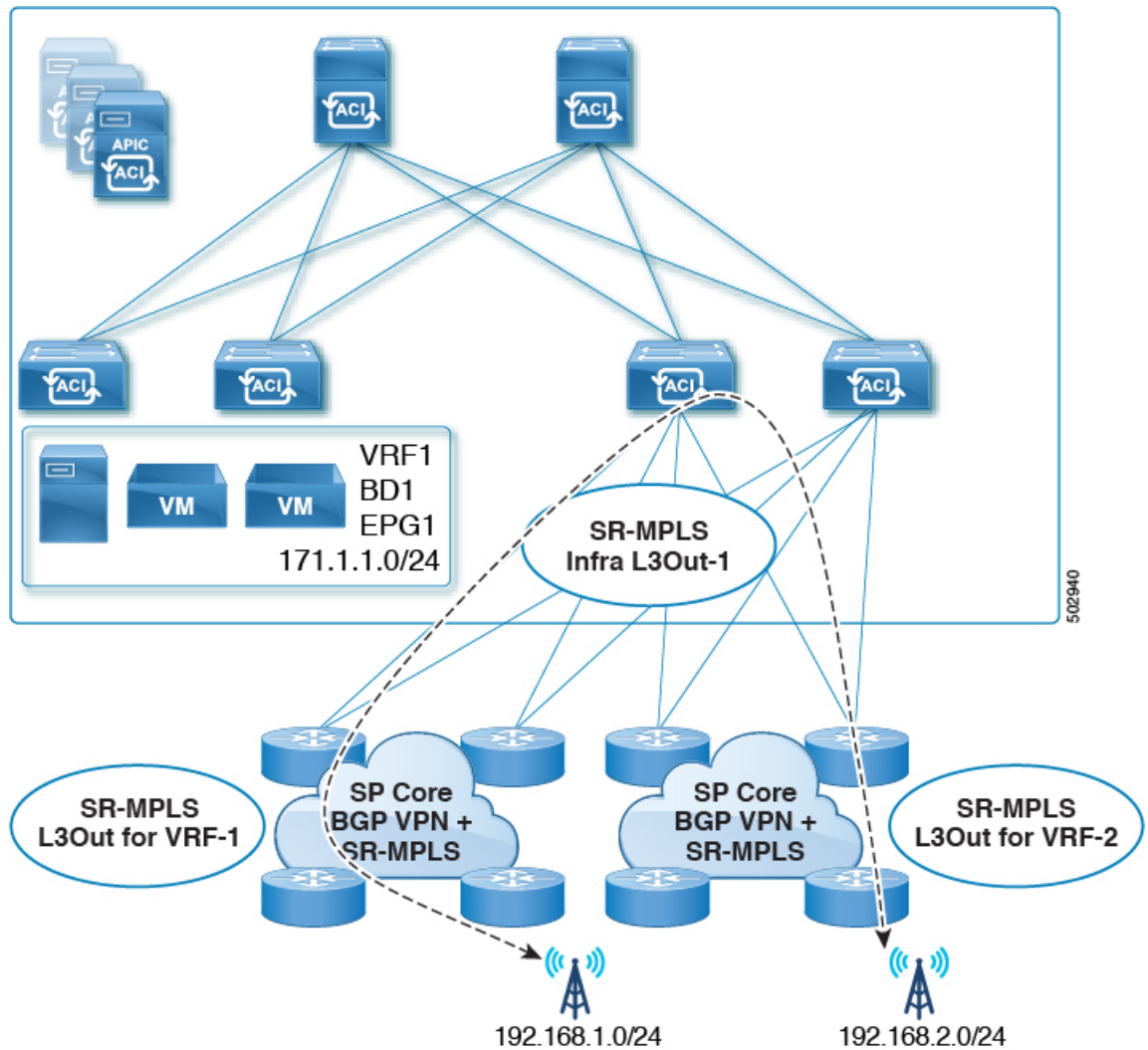


Figure 35: Transit Routing Configuration Using Different VRFs



Prior releases supported transit routing using different VRFs only.

Creating Custom QoS Policy for SR-MPLS

SR-MPLS Custom QoS policy defines the priority of the packets coming from an SR-MPLS network while they are inside the ACI fabric based on the incoming MPLS EXP values defined in the MPLS QoS ingress policy. It also marks the CoS and MPLS EXP values of the packets leaving the ACI fabric through an MPLS interface based on IPv4 DSCP values defined in MPLS QoS egress policy.



Note Creating custom QoS policy is optional. If no custom ingress policy is defined, the default QoS Level (Level3) is assigned to packets inside the fabric. If no custom egress policy is defined, the default EXP value of 0 will be marked on packets leaving the fabric.

Procedure

Step 1 Log in to your Nexus Dashboard and open the Nexus Dashboard Orchestrator service.

Step 2 Create a new Fabric Policy.

- a) From the left navigation pane, choose **Fabric Management > Fabric Policies**.
- b) On the **Fabric Policy Templates** page, click **Add Fabric Policy Template**.
- c) From the **+Create Object** dropdown, select **QoS SR-MPLS**.
- d) In the right properties sidebar, provide the **Name** for the policy.
- e) (Optional) Click **Add Description** and provide a description for the policy.

Step 3 Click **Add Ingress Rule** to add an ingress QoS translation rule.

These rules are applied for traffic that is ingressing the ACI fabric from an MPLS network and are used to map incoming packet's experimental bits (EXP) values to ACI QoS levels, as well as set DSCP and/or CoS values that the packet should have set when forwarded to an endpoint connected to the fabric.

The values are derived at the border leaf using a custom QoS translation policy. If a custom policy is not defined or not matched, default QoS Level (`Level3`) is assigned

- a) In the **Match EXP From** and **Match EXP To** fields, specify the EXP range of the ingressing MPLS packet you want to match.
- b) From the **Queuing Priority** dropdown, select the ACI QoS Level to map.

This is the QoS Level you want to assign for the traffic within ACI fabric, which ACI uses to prioritize the traffic within the fabric. The options range from `Level1` to `Level6`. The default value is `Level3`. If you do not make a selection in this field, the traffic will automatically be assigned a `Level3` priority.

- c) From the **Set DSCP** dropdown, select the DSCP value to be used when sending the un-encapsulated packet to an endpoint connected to the fabric.

The DSCP value specified is set in the original traffic received from the external network, so it will be re-exposed only when the traffic is VXLAN decapsulated on the destination ACI leaf node.

If you set the value to `Unspecified`, the original DSCP value of the packet will be retained.

- d) From the **Set CoS** dropdown, select the CoS value to be used when sending the un-encapsulated packet to an endpoint connected to the fabric.

The CoS value specified will be re-exposed only when the traffic is VXLAN decapsulated on the destination ACI leaf node.

If you set the value to `Unspecified`, the original CoS value of the packet will be retained.

In both of the above cases, the CoS preservation option must be enabled in the fabric. For more information about CoS preservation, see [Cisco APIC and QoS](#).

- e) Click the checkmark icon to save the rule.
- f) Repeat this step for any additional ingress QoS policy rules.

Step 4 Click **Add Egress Rule** to add an egress QoS translation rule.

These rules are applied for the traffic that is leaving the ACI fabric via an MPLS L3Out and are used to map the packet's IPv4 DSCP value to the MPLS packet's EXP value as well as the internal Ethernet frame's CoS value.

The setting of the packet's IPv4 DSCP value is done at the non-border leaf switch based on existing policies used for EPG and L3Out traffic. If a custom policy is not defined or not matched, the default EXP value of 0 is marked on all labels. EXP values are marked in both, default and custom policy scenarios, and are done on all MPLS labels in the packet.

Custom MPLS egress policy can override existing EPG, L3Out, and Contract QoS policies.

- a) Using the **Match DSCP From** and **Match DSCP To** dropdowns, specify the DSCP range of the ACI fabric packet you want to match for assigning the egressing MPLS packet's priority.
- b) From the **Set MPLS EXP** dropdown, select the EXP value you want to assign to the egressing MPLS packet.
- c) From the **Set CoS** dropdown, select the CoS value you want to assign to the egressing MPLS packet.
- d) Click the checkmark icon to save the rule.
- e) Repeat this step for any additional egress QoS policy rules.

- Step 5** From the **Actions** menu, select **Sites Association** and choose the SR-MPLS site with which to associate this template.
- Step 6** Click **Save** to save the template policy.
- Step 7** Click **Deploy** to deploy the fabric policy to the sites.
-

What to do next

After you have created the QoS policy, enable MPLS connectivity and configure MPLS L3Out as described in [Creating SR-MPLS Infra L3Out, on page 303](#).

Creating SR-MPLS Infra L3Out

This section describes how to configure SR-MPLS Infra L3Out settings for a site that is connected to an SR-MPLS network.

- The SR-MPLS infra L3Out is configured on the border leaf switch, which is used to set up the underlay BGP-LU and overlay MP-BGP EVPN sessions that are needed for the SR-MPLS hand-off.
- An SR-MPLS infra L3Out will be scoped to a pod or a remote leaf switch site.
- Border leaf switches or remote leaf switches in one SR-MPLS infra L3Out can connect to one or more provider edge (PE) routers in one or more routing domains.
- A pod or remote leaf switch site can have one or more SR-MPLS infra L3Outs.

Before you begin

You must have:

- Added a site that is connected via SR-MPLS network as described in [Adding Cisco ACI Sites, on page 143](#).
- If necessary, created SR-MPLS QoS policy as described in [Creating Custom QoS Policy for SR-MPLS, on page 301](#).

Procedure

Step 1 Ensure that SR-MPLS Connectivity is enabled for the site.

- a) In the main navigation menu, select **Infrastructure > Site Connectivity**.
- b) In the **Site Connectivity** page, click **Configure**.
- c) In the left pane, under **Sites**, select the specific site that is connected via SR-MPLS.
- d) In the right **<Site> Settings** pane, enable the **SR-MPLS Connectivity** and provide the SR-MPLS information.
 - The **Segment Routing Global Block (SRGB) Range** is the range of label values reserved for Segment Routing (SR) in the Label Switching Database (LSD). The Segment ID (SID) is a unique identifier for a specific segment and is configured on each node for the MPLS transport loopback. The SID value, which you will configure later as part of the border leaf configuration, is advertised using BGP-LU to the peer router, and the peer router uses the SID index to calculate the local label.

The default range is 16000–23999.

- The **Domain ID Base** enables the BGP Domain-Path feature. For more information, see [Cisco APIC Layer 3 Networking Configuration Guide](#).

If you choose to provide a value in this field to enable the Domain-Path feature, ensure that you use a unique value for each SR-MPLS site in your Multi-Site domain, which will be specific to this ACI fabric.

Step 2 In the main pane, click **+Add SR-MPLS L3Out** within a pod.

Step 3 In the right **Properties** pane, provide a name for the SR-MPLS L3Out.

Step 4 (Optional) From the **QoS Policy** dropdown, select a QoS Policy you created for SR-MPLS traffic.

Select the QoS policy you created in [Creating Custom QoS Policy for SR-MPLS, on page 301](#).

Otherwise, if you do not assign a custom QoS policy, the following default values are assigned:

- All incoming MPLS traffic on the border leaf switch is classified into QoS Level 3 (the default QoS level).
- The border leaf switch does the following:
 - Retains the original DSCP values for traffic coming from SR-MPLS without any remarking.
 - Forwards packets to the MPLS network with the original CoS value of the tenant traffic if the CoS preservation is enabled.
 - Forwards packets with the default MPLS EXP value (0) to the SR-MPLS network.
- In addition, the border leaf switch does not change the original DSCP values of the tenant traffic coming from the application server while forwarding to the SR network.

Step 5 From the **L3 Domain** dropdown, select the Layer 3 domain.

Step 6 Configure settings for border leaf switches and ports connected to the SR-MPLS network.

You need to provide information about the border leaf switches as well as the interface ports which connect to the SR-MPLS network.

- a) Click **+Add Leaf** to add a leaf switch.
- b) In the **Add Leaf** window, select the leaf switch from the **Leaf Name** dropdown.

- c) In the **SID Index** field, provide a valid segment ID (SID) offset.

When configuring the interface ports later in this section, you will be able to choose whether you want to enable segment routing. The SID index is configured on each node for the MPLS transport loopback. The SID index value is advertised using BGP-LU to the peer router, and the peer router uses the SID index to calculate the local label. If you plan to enable segment routing, you must specify the segment ID for this border leaf.

- The value must be within the SRGB range you configured earlier.
- The value must be the same for the selected leaf switch across all SR-MPLS L3Outs in the site.
- The same value cannot be used for more than one leaf across all sites.
- If you need to update the value, you must first delete it from all SR-MPLS L3Outs in the leaf and re-deploy the configuration. Then you can update it with the new value, followed by re-deploying the new configuration.

- d) Provide the local **Router ID**.

Unique router identifier within the fabric.

- e) Provide the **BGP EVPN Loopback** address.

The BGP-EVPN loopback is used for the BGP-EVPN control plane session. Use this field to configure the MP-BGP EVPN session between the EVPN loopbacks of the border leaf switch and the DC-PE to advertise the overlay prefixes. The MP-BGP EVPN sessions are established between the BGP-EVPN loopback and the BGP-EVPN remote peer address (configured in the **MPLS BGP-EVPN Peer IPv4 Address** field in the **BGP Connectivity** step before).

While you can use a different IP address for the BGP-EVPN loopback and the MPLS transport loopback, we recommend that you use the same loopback for the BGP-EVPN and the MPLS transport loopback on the ACI border leaf switch.

- f) Provide the **MPLS Transport Loopback** address.

The MPLS transport loopback is used to build the data plane session between the ACI border leaf switch and the DC-PE, where the MPLS transport loopback becomes the next-hop for the prefixes advertised from the border leaf switches to the DC-PE routers.

While you can use a different IP address for the BGP-EVPN loopback and the MPLS transport loopback, we recommend that you use the same loopback for the BGP-EVPN and the MPLS transport loopback on the ACI border leaf switch.

- g) Click **Add Interface** to provide switch interface details.

From the **Interface Type** dropdown, select whether it is a Layer 3 physical interface or a port channel interface. If you choose to use a port channel interface, it must have been already created on the APIC.

Then provide the interface, its IP address, and MTU size. If you want to use a subinterface, provide the **VLAN ID** for the sub-interface, otherwise leave the VLAN ID field blank.

In the **BGP-Label Unicast Peer IPv4 Address** and **BGP-Label Unicast Remote AS Number**, specify the BGP-LU peer information of the next hop device, which is the device connected directly to the interface. The next hop address must be part of the subnet configured for the interface.

Choose whether you want to enable an MPLS or an SR-MPLS hand-off.

(Optional) Choose to enable the additional BGP options based on your deployment.

Finally, click the checkmark to the right of the **Interface Type** dropdown to save interface port information.

- h) Repeat the previous sub-step for all interfaces on the switch that connect to the MPLS network.

- i) Click **Save** to save the leaf switch information.
- j) Repeat this step for all leaf switches connected to MPLS networks.

Step 7 Configure BGP settings.

You must provide BGP connectivity details for the BGP EVPN connection between the site's border leaf (BL) switches and the provider edge (PE) router.

- a) Click **+Add BGP-EVPN Connectivity**.
- b) In the **Add MPLS BGP-EVPN Connectivity** window, provide the details.

For the **MPLS BGP-EVPN Peer IPv4 Address** field, provide the loopback IP address of the DC-PE router, which is not necessarily the device connected directly to the border leaf.

For the **Remote AS Number**, enter a number that uniquely identifies the neighbor autonomous system of the DC-PE. the Autonomous System Number can be in 4-byte as plain format from 1 to 4294967295. Keep in mind, ACI supports only `asplain` format and not `asdot` or `asdot+` format AS numbers. For more information on ASN formats, see [Explaining 4-Byte Autonomous System \(AS\) ASPLAIN and ASDOT Notation for Cisco IOS](#) document.

For the **TTL** field, specify a number large enough to account for multiple hops between the border leaf and the DC-PE router, for example 10. The allowed range is 2–255 hops.

(Optional) Choose to enable the additional BGP options based on your deployment.

- c) Click **Save** to save BGP settings.
- d) Repeat this step to for any additional BGP connections.

Typically, you would be connecting to two DC-PE routers, so provide BGP peer information for both connections.

Step 8 Deploy the changes to sites.

What to do next

After you have enabled and configured MPLS connectivity, you can create and manage Tenants, route maps, and schemas as described in the [Creating SR-MPLS Route Map Policy, on page 306](#).

Creating SR-MPLS Route Map Policy

This section describes how to create a route map policy. Route maps are sets of `if-then` rules that enable you to specify which routes are advertised out of the Tenant SR-MPLS L3Out. Route maps also enable you to specify which routes received from the DC-PE routers will be injected into the BGP VPNv4 ACI control plane.

You will use the SR-MPLS route map policy in the next section when defining site-local settings for the SR-MPLS L3Out.

Procedure

Step 1 Log in to your Nexus Dashboard and open the Nexus Dashboard Orchestrator service.

Step 2 Create a new Tenant Policy.

- a) From the left navigation pane, choose **Application Management > Tenant Policies**.

- b) On the **Tenant Policy Templates** page, click **Add Tenant Policy Template**.
- c) In the Tenant Policies page's right properties sidebar, provide the **Name** for the tenant.
- d) From the **Select a Tenant** dropdown, choose the tenant with which you want to associate this template.

All the policies you create in this template as described in the following steps will be associated with the selected tenant deployed to it when you push the template to a specific site.

By default, the new template is empty, so you need to add one or more tenant policies as described in the following steps. Note that you don't have to create every policy available in the template – you can create a template with just a single route map policy for your SR-MPLS use case.

Step 3

Create a Route Map Policy for Route Control.

- a) From the **+Create Object** dropdown, select **Route Map Policy for Route Control**.
- b) In the right properties sidebar, provide the **Name** for the policy.
- c) (Optional) Click **Add Description** and provide a description for the policy.
- d) Click **+Add Entry** and provide the route map information.

For each route map, you need to create one or more context entries. Each entry is a rule that defines an action based on one or more matching criteria based on the following information:

- **Context Order** – Context order is used to determine the order in which contexts are evaluated. The value must be in the 0–9 range.
- **Context Action** – Context action defines the action to perform (`permit` or `deny`) if a match is found.

Once the context order and action are defined, choose how you want to match the context:

- Click **+Add Attribute** to specify the action that will be taken should the context match.

You can choose one of the following actions:

- Set Community
- Set Route Tag
- Set Dampening
- Set Weight
- Set Next Hop
- Set Preference
- Set Metric
- Set Metric Type
- Set AS Path
- Set Additional Community

After you have configured the attribute, click **Save**.

- If you want to match an action based on an IP address or prefix, click **Add IP Address**.

In the **Prefix** field, provide the IP address prefix. Both IPv4 and IPv6 prefixes are supported, for example `2003:1:1a5:1a5::/64` or `205.205.0.0/16`.

If you want to aggregate IPs in a specific range, check the **Aggregate** checkbox and provide the range. For example, you can specify `0.0.0.0/0` prefix to match any IP or you can specify `10.0.0.0/8` prefix to match any `10.x.x.x` addresses.

- If you want to match an action based on community lists, click **Add Community**.

In the **Community** field, provide the community string. For example, `regular:as2-nn2:200:300`.

Then choose the **Scope**: `Transitive` means the community will be propagated across eBGP peering (across autonomous systems) while `Non-Transitive` means the community will not be propagated.

- e) Repeat the previous substeps to create any additional route map entries for the same policy.
- f) Click **Save** to save the policy and return to the template page.
- g) Repeat this step to create any additional Route Map for Route Control policies.

Step 4 From the **Actions** menu, select **Sites Association** and choose the SR-MPLS site with which to associate this template.

Step 5 Click **Deploy** to deploy the tenant policy to the sites.

Configure EPG-to-External-EPG (North-South) Communication

This section describes how to establish North-South communication between an application EPG and an external SR-MPLS network. You can also use this approach to enable EPG-to-EPG communication across sites via the SR-MPLS L3Out data path (leveraging the external SR-MPLS network).

If instead you want to establish EPG-to-EPG intersite connectivity via the VXLAN data plane across the ISN which is supported starting with release 4.0(2), you can simply establish a contract relationship between those EPGs as you typically would.

Procedure

Step 1 Choose the template or create a new one.

You can select the template as you typically would for other ACI fabric use cases:

- a) In the main navigation menu, select **Application Management > Schemas**.
- b) Select an existing schema or create a new one.
- c) Select an existing template or click **Add New Template** and select `ACI Multi-Cloud` for template type.
- d) Select the tenant for the new template.

The tenant must be associated with the SR-MPLS site.

- e) (Optional) Enable the **Autonomous** option for the template if you plan to deploy this template only to sites that do not have any intersite connectivity to other sites.

Step 2 Create a VRF.

- a) From the **+Create Object** menu, choose **VRF**.
- b) In the right properties sidebar, provide the name for the VRF.

Step 3 Create an SR-External EPG.

Note

If you assign the template that contains SR-External EPG to multiple sites, the EPG will be stretched to all of those sites. In this case each site must have a local SR-MPLS L3Out or you will not be allowed to deploy that template to all associated sites.

- a) From the **+Create Object** menu, choose **SR-External EPG**.
- b) In the right properties sidebar, provide the name for the external EPG.
- c) From the **Virtual Routing & Forwarding** dropdown, select the VRF you created in the previous step.

Step 4

Create an SR-MPLS L3Out.

Note

The tenant SR-MPLS L3Out must be defined in the same template as the SR-External EPG you created in the previous step.

- a) From the **+Create Object** menu, choose **SR-L3Out**.
- b) In the right properties sidebar, provide the name for the L3Out.
- c) From the **Virtual Routing & Forwarding** dropdown, select the same VRF you selected for the external EPG in the previous step.
- d) From the **SR-External EPGs** dropdown, select the SR-External EPG you created in the previous step.

Step 5

Assign the template to a single site or to multiple sites, depending on the specific use case you need to configure.

Step 6

Select the site-local settings for the template you are configuring.

In the following few steps, you will configure site local settings for the VRF, SR-External EPG, and SR-MPLS L3Out you created in the previous steps.

Step 7

Configure site-local settings for the VRF.

You must provide BGP route information for the VRF used by the SR-MPLS L3Out.

- a) In the main pane, scroll down to **VRF** area and select the VRF you created in the previous step.
- b) From the **Address Family** dropdown, select whether it is IPv4 or IPv6 address.
- c) In the **Route Target** field, provide the route string.

For example, `route-target:ipv4-nn2:1.1.1.1:1901`.

- d) From the **Type** dropdown, select whether to import or export the route.
- e) Click **Save** to save the route information.
- f) (Optional) Repeat this step to add any additional BGP route targets.

Step 8

Configure site-local settings for the SR-MPLS L3Out.

- a) In the main pane, scroll down to **SR-MPLS L3Out** area and select the MPLS L3Out.
- b) Click **+Add SR-MPLS Infra L3Out** to add the Infra SR-MPLS L3Out information.
- c) In the **Add SR-MPLS Infra L3Out** window, select the Infra SR-MPLS L3Out you created when configuring Infra for that site.
- d) Click **Add Route Map Policy** and select the route map policy you created in the tenant policy template in previous section, then specify whether you want to import or export the routes.

You must configure a single export route map policy. Optionally, you can configure an additional import route map policy.

- e) Click **Ok** to save the changes.

Step 9

Create and configure an application EPG as you typically would.

Note

The EPG can be in the same or different template and schema.

Step 10 Create a contract between the application EPG and the SR-External EPG.

Step 11 Deploy the configuration

- a) In the main pain of the **Schemas** view, click **Deploy to Sites**.
- b) In the **Deploy to Sites** window, verify the changes that will be pushed to the site and click **Deploy**.

Note

Starting from release 4.0(2), it is possible to use the EPG-to-SR-External-EPG contracts exclusively for North-South traffic (communication with resources external to the ACI fabrics), similar to the traditional IP-based L3Outs. In that case, EPG-to-EPG intersite communication can be enabled via the VXLAN data path across the ISN by simply creating a contract relationship between those EPGs.

However, if you want to establish EPG-to-EPG (East-West) communication between EPGs in different sites across the external SR-MPLS network, you can do that as outlined in the next step.

Step 12 If you want to use the SR-MPLS L3Out data path for EPG-to-EPG traffic across sites (leveraging the SR-MPLS external network instead of the VXLAN data path across the ISN), you can establish contracts between each site-local EPG and the SR-External EPG associated to the tenant SR-MPLS L3Out.

The SR-External EPG can be deployed as a site-local object in each site or as a stretched object across sites. Note that using the SR-MPLS L3out data path for EPG-to-EPG traffic across sites is only possible if there are no direct contract relationships between those EPGs or between each EPG and any other remote EPG

- a) Create two application EPGs as you typically would in templates associated to different sites.

For example, `epg1` and `epg2`.

These EPGs can be in the same or different VRFs or Tenants.

- b) Create two separate site-local SR-External EPGs or a single stretched SR-External EPG.

If you are creating separate SR-External EPGs, they can be in the same or different VRFs or Tenants and the same template or different templates depending on the specific deployment scenario.

For example, `mpls-extepg-1` and `mpls-extepg-2`

- c) Configure two separate site-local tenant SR-MPLS L3Outs or a single stretched tenant SR-MPLS L3Out.

For example, `mpls-l3out-1` and `mpls-l3out-2`

- d) Create a contract you will use to allow traffic between each site local EPG and SR-MPLS L3Out local connection.

You will need to create and define a filter for the contract as you typically would.

- e) Assign the contracts to the appropriate EPGs.

In order to allow traffic between the two application EPGs you created, you will actually need to assign the contract twice: once between `epg1` and its `mpls-extepg-1` and then again between `epg2` and its `mpls-extepg-2`. Note that it's possible to have the same SR-External EPG instead of two separate ones if it is stretched across sites.

As an example, if you want `epg1` to provide a service to `epg2`, you would:

- Assign the contract to `epg1` with type `provider`.
- Assign the contract to `mpls-l3out-1` with type `consumer`.
- Assign the contract to `epg2` with type `consumer`.

- Assign the contract to `mpls-extepg-2` with type `provider`.
-



CHAPTER 26

vzAny Contracts

- [vzAny and Multi-Site, on page 313](#)
- [vzAny and Multi-Site Guidelines and Limitations, on page 314](#)
- [Create Contract and Filters, on page 316](#)
- [Configure vzAny to Consume/Provide a Contract, on page 317](#)
- [Create EPGs to Be Part of the vzAny VRF, on page 317](#)
- [Free Intra-VRF Communication, on page 318](#)
- [Many-to-One Communication, on page 323](#)

vzAny and Multi-Site

The `vzAny` managed object provides a convenient way of associating all endpoint groups (EPGs) in a Virtual Routing and Forwarding (VRF) instance to one or more contracts, instead of creating a separate contract relation for each EPG.

In the Cisco ACI fabric, EPGs can only communicate with other EPGs according to contract rules. A relationship between an EPG and a contract specifies whether the EPG provides the communications defined by the contract rules, consumes them, or both. By dynamically applying contract rules to all EPGs in a VRF, `vzAny` automates the process of configuring EPG contract relationships. Whenever a new EPG is added to a VRF, `vzAny` contract rules automatically apply. The `vzAny` one-to-all EPG relationship is the most efficient way of applying contract rules to all EPGs in a VRF.



Note External EPGs that are associated with L3Outs and are part of a VRF are also included in the `vzAny` logical group.

Advantages

Policy information in Cisco ACI is programmed in the fabric switches' TCAM tables. TCAM entries are typically specific to each pair of EPGs that are allowed to communicate with each other via a Contract. This means that even if the same contract is re-used, multiple TCAM entries are created for every pair of EPGs.

The size of the policy TCAM table depends on the generation of the switches that you are using. In certain large scale environments it is important to take policy TCAM usage into account and ensure that the limits are not exceeded.

vzAny allows you to combine all EPGs within the same VRF into a single "group" and create a contract relationship with that group rather than individual EPGs within it, while consuming only a single TCAM entry. This saves the time you would otherwise spend creating multiple contract relationships for individual EPGs in the VRF as well as the TCAM space.

Use Cases

There are two typical use cases for vzAny:

- Free communication between EPGs within the same VRF, as described in [Free Intra-VRF Communication, on page 318](#).
- Many-to-one communication allowing all EPGs within the same VRF to consume a shared service from a single EPG, as described in more detail in [Many-to-One Communication, on page 323](#).

vzAny and Multi-Site Guidelines and Limitations

The following guidelines and limitations apply when using vzAny:

- If you plan to enable the vzAny object for a given VRF to provide or consume a contract, the following additional restrictions apply:
 - If vzAny for a given VRF is configured as consumer of a contract c1, the vzAny objects for other VRFs must not be configured as providers of c1.
 - If vzAny for a given VRF is configured as provider of a contract c1, the vzAny objects for other VRFs must not be configured as consumers of c1.
 - If an External EPG part of a given VRF is consuming a contract c1, the vzAny objects for other VRFs must not be configured as providers of c1.
 - If an EPG part of a given VRF is consuming a contract c1, the vzAny objects for other VRFs must not be configured as providers of c1.
 - If vzAny for a given VRF is configured as provider of a contract c1, then EPGs, External EPGs or vzAny objects for other VRFs must not be configured as consumers of c1.
- EPGs and External EPGs objects in a given VRF must not be configured as part of the Preferred Group if vzAny for that VRF is already consuming or providing a contract.
- If any EPG or External EPG objects in a given VRF are deployed in a cloud site, it is not possible to configure vzAny for that VRF to consume or provide a contract
- vzAny is supported with inter-VRF intersite L3Out configurations only when the fabrics are part of the Multi-Site domain running Cisco ACI 5.2(4) release or later.
- vzAny must not consume or provide a contract that is associated with a Service-Graph with PBR.
- vzAny can be configured as provider, consumer or both of a contract for establishing intra-VRF communication.
- vzAny is supported only as a consumer of a shared service but not as a provider.
- We recommend stretching the vzAny VRF to all sites where you plan to deploy EPGs and BDs that use it.

- You can import existing vzAny configurations from an APIC.

**Note**

In certain cases due to an existing issue ([CSCvt47568](#)), if you make changes to the imported configuration before re-deploying it from the Nexus Dashboard Orchestrator, some changes may not get correctly updated in the APICs. To avoid this, re-deploy the configuration immediately after importing but before making any changes to it. After you re-deploy the unchanged config, you will be able to update it as normal.

- vzAny providers and consumers include application EPGs, external EPGs associated to L3Outs, and endpoint groups for in-band or out-of-band access.
- vzAny implicitly creates a 0.0.0.0/0 classification for externally originating traffic, allowing all traffic originating from any external IP subnet. When vzAny is in use for a VRF, it also includes the external EPGs associated to the L3Outs part of that VRF, hence it is equivalent to having created a L3external classification that includes the subnets specified in the VRF itself.
- If an EPG within a VRF is consuming a shared service contract provided by an EPG in a different VRF, the traffic from the EPG of the provider VRF is filtered within the consumer VRF. vzAny is equivalent to a wildcard for the source or destination EPG.

Be careful when you configure a shared service contract between vzAny in the consumer VRF and an EPG1 in a different provider VRF. Since the policy enforcement (filtering) is always performed in the consumer VRF, if the subnet associated to another EPG2 that is part of the provider VRF is leaked into the consumer VRF, then EPG2 will start communicating with consumer EPGs across VRFs even without explicitly providing a contract. Failure to observe this guideline could allow unintended traffic between EPGs across VRFs.

- Configuring a VRF with vzAny as both provider and consumer of a contract using an "allow all" filter, is the same as configuring an unenforced VRF. This implies that all EPGs within that VRF are free to communicate to each other without a contract.
- If the contract scope is application-profile, the vzAny configuration is ignored and filter rules are expanded; CAM utilization is the same as if specific contracts were deployed between each pair of consumer and provider EPGs. In this case, there is no benefit in terms of TCAM space usage.
- In the case of shared services, you must define the provider EPG shared subnet under the EPG in order to properly derive the classification (`pcTag`) of the destination on the consumer (vzAny) side. If you are migrating from a BD-to-BD shared services configuration, where both the consumer and provider subnets are defined under bridge domains, to vzAny acting as a shared service consumer, you must take an extra configuration step where you add the provider subnet to the EPG with the shared flags at minimum. However, since the subnet under the EPG is not needed for connectivity, it is always recommended to check the `No default SVI gateway` flag.

If you add the EPG subnet as a duplicate of the defined BD subnet, ensure that both definitions of the subnet always have the same flags defined. Failure to do so can result in an error.

Create Contract and Filters

When using vzAny, you are essentially creating a single point for a contract relationship, as such you must have a typical contract you will use for any such relationship as well as the filter for the contract.

This section describes how to create a new contract specifically for this purpose. Alternatively, you can choose to import any existing contracts you have configured on each APIC site.

Procedure

Step 1 Log in to the Nexus Dashboard Orchestrator GUI.

Step 2 From the left navigation pane, select **Schemas**.

Step 3 Select the Schema where you want to create your Contract.

If you have an existing Schema you want to update, simply click the Schema's name in the main window pane. Otherwise, if you want to create a new Schema, click the **Add Schema** button and provide the schema information, such as the name and tenant, as you typically would.

Step 4 Create a filter.

- a) Scrolls down to the **Filter** area and click the + sign to add a new filter.
- b) Provide the name for the Contract.
- c) Click **+Entry** to add a filter entry.
- d) In the **Add Entry** window, provide filter details.

Provide the filter details as you typically would to define the kind of traffic you want to allow.

- e) Click **SAVE** to add the entry.
- f) (Optional) If required, create additional filter entries.

Step 5 Create the contract.

- a) Scrolls down to the **Contract** area and click the + sign to add a new contract.
- b) Provide the name for the Contract.

For example, `contract-vzany`.

- c) Choose the scope for the contract

Choose the scope appropriate for your use-case. For example, if you want to enable cross-tenant shared services, you must set the scope to `Global`.

- d) Choose whether the contract will apply in both directions
- e) Click **+Filter** to add one or more contract filters.
- f) In the **Add Filter Chain** window, choose the filter you created in the previous step.
- g) Click **SAVE** to add the filter.
- h) (Optional) If required, repeat the procedure to provide additional filters.
- i) (Optional) If you disabled the **Apply Both Directions** option, provide filters for both, consumer and provider directions.

You have now created the contract you will use with vzAny in the next section.

Configure vzAny to Consume/Provide a Contract

This section describes how to create a vzAny VRF or enable an existing VRF for vzAny.

Before you begin

You must have:

- Created a Contract and one or more Filters to use with vzAny as described in [Create Contract and Filters, on page 316](#).

Procedure

-
- | | |
|---------------|---|
| Step 1 | Log in to the Nexus Dashboard Orchestrator GUI. |
| Step 2 | From the left navigation pane, select Schemas . |
| Step 3 | Select the Schema containing the specific template with the definition of the VRF.

For new configuration, you can create a new schema with the Add Schema button and then define a new template (associated to the tenant of interest) where you can configure VRF. |
| Step 4 | Create or select a VRF.

If you have an existing VRF for which you want to configure vzAny to provide/consume a contract, simply click the VRF in the main window pane. Otherwise, if you want to create a new VRF, scroll down to the VRF area and click the + sign. |
| Step 5 | Select vzAny.

In the right sidebar, check the vzAny checkbox. |
| Step 6 | Select the vzAny contract.

The +Contract option becomes available after you enable the vzAny checkbox.

a) Click +Contract to add the contract
b) Select the contract.

Select the contract you created in Create Contract and Filters, on page 316 .

c) Select the Contract type.

You can choose either <code>consumer</code> or <code>provider</code> for the contract based on your use case. |
-

Create EPGs to Be Part of the vzAny VRF

You can choose to create new or use existing EPGs for your vzAny use cases. There are no explicit vzAny settings on the EPGs and as soon as an EPG is associated to a BD in a VRF, the EPG becomes part of the

vzAny logical group for that VRF (the vzAny VRF). If you simply enabled vzAny for an existing VRF with all its EPGs already created and configured, you can skip this section.

Before you begin

You must have:

- Created a Contract and one or more Filters to use with vzAny as described in [Create Contract and Filters, on page 316](#).
- Created the vzAny VRF and assigned the Contract to it as described in [Configure vzAny to Consume/Provide a Contract, on page 317](#).

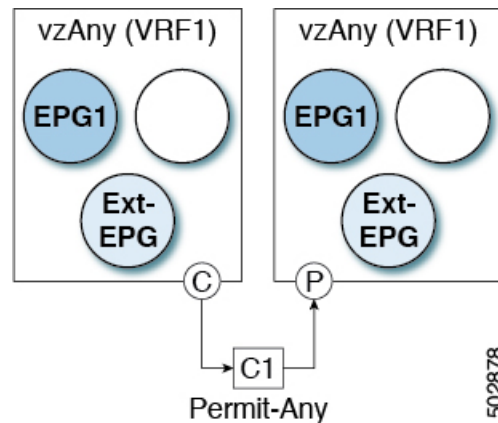
Procedure

-
- Step 1** If you want to create an EPG to be part of the vzAny VRF
- a) Create a BD you will use for your EPG.
 - b) In the BD configuration sidebar's **Virtual Routing & Forwarding** dropdown, select the vzAny VRF you created.
 - c) Create an EPG.
 - d) In the EPG configuration sidebar's **Bridge Domain** dropdown, select the BD you created.
- Step 2** If you want to create an External EPG to be part of the vzAny VRF
- a) Create an external EPG.
 - b) In the External EPG configuration sidebar's **Virtual Routing & Forwarding** dropdown, select the vzAny VRF you created.
-

Free Intra-VRF Communication

This section shows a number of schema examples for unrestricted intra-VRF communication. In all shown scenarios vzAny provides and consumes a contract with a `permit-any` filter. This essentially uses the ACI fabrics for network connectivity only without any policy enforcement and is equivalent to the *VRF Unenforced* option.

Figure 36:



For all the following use cases, you will need to create the same objects and policies summarized below. However, the schema and template design will depend on the number of sites as well as which objects are going to be stretched. The specific sections below contain recommendation on template layout.

Procedure

-
- Step 1** Create a Schema.
- Step 2** Create a common Template used to deploy configuration objects in all the sites (that is *stretched objects*).
- Step 3** Create any additional templates for every combination of sites where EPGs will be deployed .
- If you will deploy a single template to all sites, you can skip this step. The use-case diagrams in the following sections provide template examples.
- Step 4** Within the common Template, create the contract and filters to be consumed/provided by vzAny.
- In this specific use case, the contract should have a single "permit-any" filter rule.
- For specific steps, see [Create Contract and Filters, on page 316](#).
- Step 5** Within the common Template, create a VRF and configure vzAny to consume and provide the previously defined contract with the "permit-any" rule.
- This ensures that free intra-VRF communication can be established.
- For specific steps, see [Configure vzAny to Consume/Provide a Contract, on page 317](#).
- Step 6** Within each site's template, create and configure the EPGs that will be deployed to that site only.
- If you will deploy a single template to all sites, create the EPGs within the same template as the VRF instead. The use-case diagrams in the following sections provide template examples.
- This is described in [Create EPGs to Be Part of the vzAny VRF, on page 317](#).
- Step 7** Assign the common Template to every site.
- Step 8** Assign each template to the appropriate sites.
- Step 9** Deploy the templates.
-

Stretched EPGs

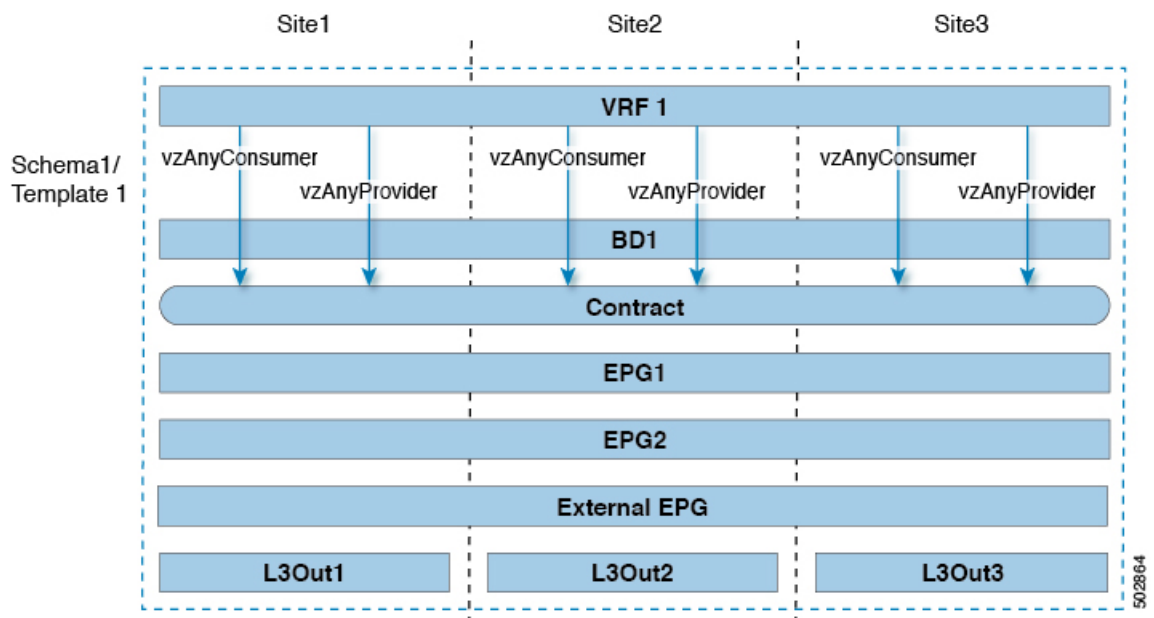
The following example shows intra-VRF communication between EPGs or External EPGs all of which are stretched between sites. In this example EPG1 and EPG2 are mapped to the same BD1, but they could each be part of different BDs as long as both BDs are part of VRF1.

In this case you can create all objects within the same template and then deploy the template to all sites.



Note As a best practice, it is recommended that the L3Out objects should instead remain defined only on Cisco APIC or configured on-site local templates on MSO.

Figure 37:



Site-Local EPGs

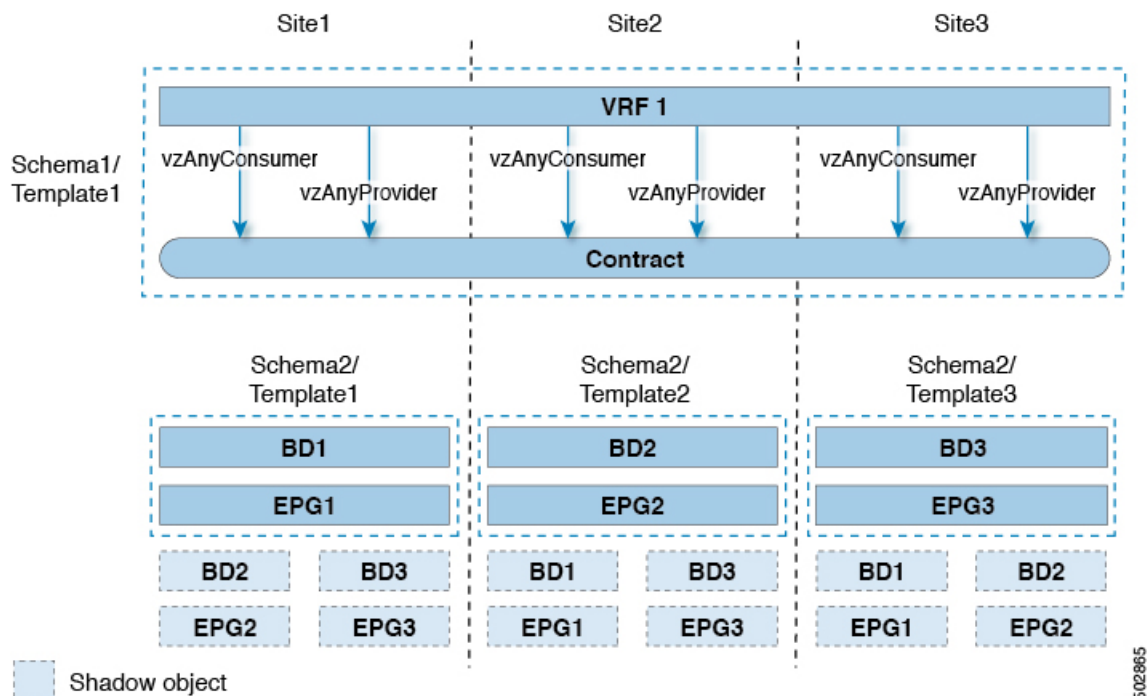
The following example shows intra-VRF communication between EPGs or External EPGs where none of the EPGs are stretched but can still freely communicate with each other since vzAny consumes and provides the "permit-any" contract.

In this case you will need to create multiple templates:

- A single template for the shared objects (VRF, Contract) deployed to every site.
- And a separate template for every site containing the EPG and BD deployed that site.

For the objects that are not stretched, shadow objects are created in other sites.

Figure 38:



502865

Combination of Site-Local and Stretched EPGs

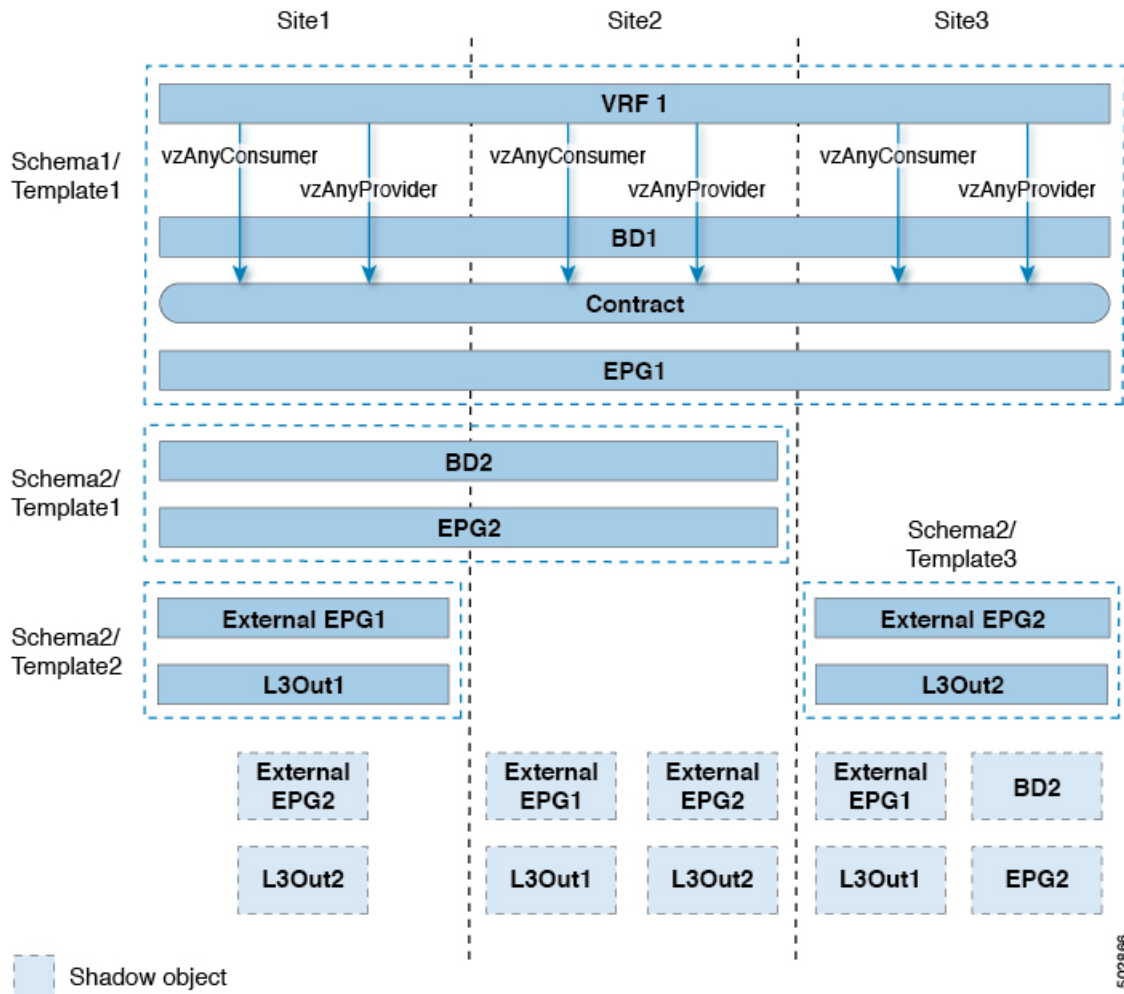
The following example shows intra-VRF communication between EPGs or External EPGs where some EPGs are stretched while others are deployed to a single site only. All EPGs can still freely communicate with each other since vzAny consumes and provides the "permit-any" contract.

In this case you will need to create multiple templates:

- A single template for the shared objects (VRF, Contract, BDs) deployed to every site.
- And a separate template for every site combination containing the objects deployed only to those sites.

For the objects that are not stretched, shadow objects are created in other sites.

Figure 39:



502866

Intra-VRF Intersite L3Out

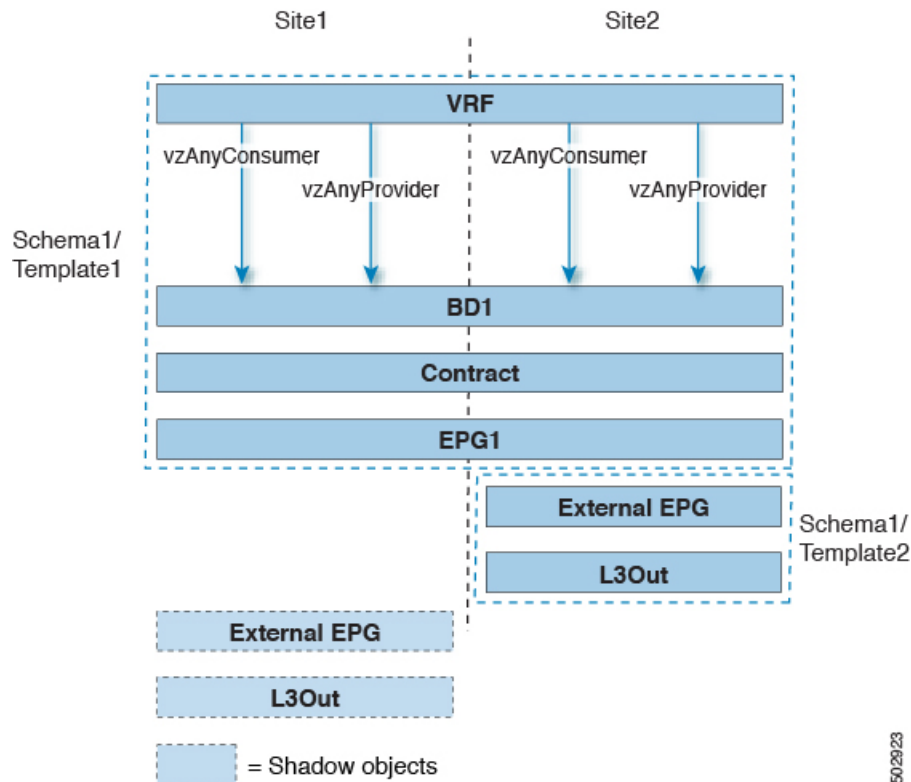
This use case allows you to configure an intersite L3Out for multiple EPGs within a vzAny VRF. When the L3Out's external EPG is in the same VRF, you do not need to explicitly add the provider contract to the external EPG.

Keep in mind, when configuring an intersite L3Out, you must configure a routable TEP pool for each Pod. Additional intersite L3Out details and requirements are described in the [Intersite L3Out Overview](#), on page 205 section.

In this case you will need to create multiple templates:

- A single template for the shared vzAny objects (VRF, Contract, BD) deployed to one or more sites.
- And a separate template for every site combination containing the objects deployed only to those sites.

Figure 40:



Based on the configuration shown in the above figure, endpoints that are part of the stretched EPG1 and connected to Site1 will be able to communicate with the external network domain via the L3Out connection deployed in Site2. The same would apply for endpoints that are part of site-local EPGs deployed in Site1.

Inter-VRF Intersite L3Out

This use case allows you to enable vzAny contracts between a consumer VRF and L3Out external EPGs in a different provider VRF. Multiple EPGs that are part of the vzAny consumer VRF can communicate with a single EPG that is providing a shared service in a provide VRF. vzAny contract acts as a contract for all EPGs in a VRF. Each of the participating VRFs and L3Out external EPG can be stretched across sites.



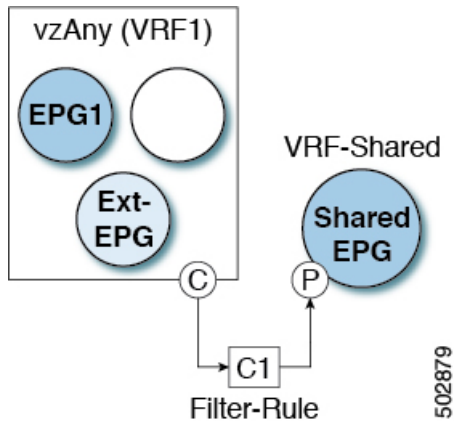
Note A VRF cannot be a vzAny provider.

Many-to-One Communication

The following three sections provide schema examples of multiple EPGs that are part of the same vzAny VRF communicating with a single EPG that is providing a shared service. In this case, the contract can specify one or more filter rules.

The EPG providing shared services can be in a separate VRF (as shown in the figure below) or it can be part of the vzAny VRF.

Figure 41:



For all the following use cases, you will need to create the same objects and policies summarized below. However, the schema and template design will depend on the number of sites as well as which objects are going to be stretched. The specific sections below contain recommendation on template layout.

Procedure

-
- Step 1** Create a Schema.
- Step 2** Create a common Template used to deploy configuration objects in all the sites (that is *stretched objects*).
- Step 3** Create any additional templates for every combination of sites where EPGs will be deployed .
- Step 4** Within the common Template, create the contract and filters to be consumed by vzAny and provided by the EPG offering shared services.
- This is described in [Create Contract and Filters, on page 316](#).
- Step 5** Within the common Template, create a VRF and configure vzAny to consume the previously defined contract.
- This is described in [Configure vzAny to Consume/Provide a Contract, on page 317](#).
- Step 6** Within each site's template, create and configure the EPGs that are part of the vzAny VRF.
- This is described in [Create EPGs to Be Part of the vzAny VRF, on page 317](#).
- Step 7** Create new or configure existing provider EPG or external EPG.
- You create and configure the provider EPG or external EPG as you typically would.
- Step 8** Assign the Contract to the provider EPG.
- In addition to assigning the contract to be consumed by vzAny, you will also need to assign the same contract to the provider EPG.
-

Provider EPG Within vzAny VRF

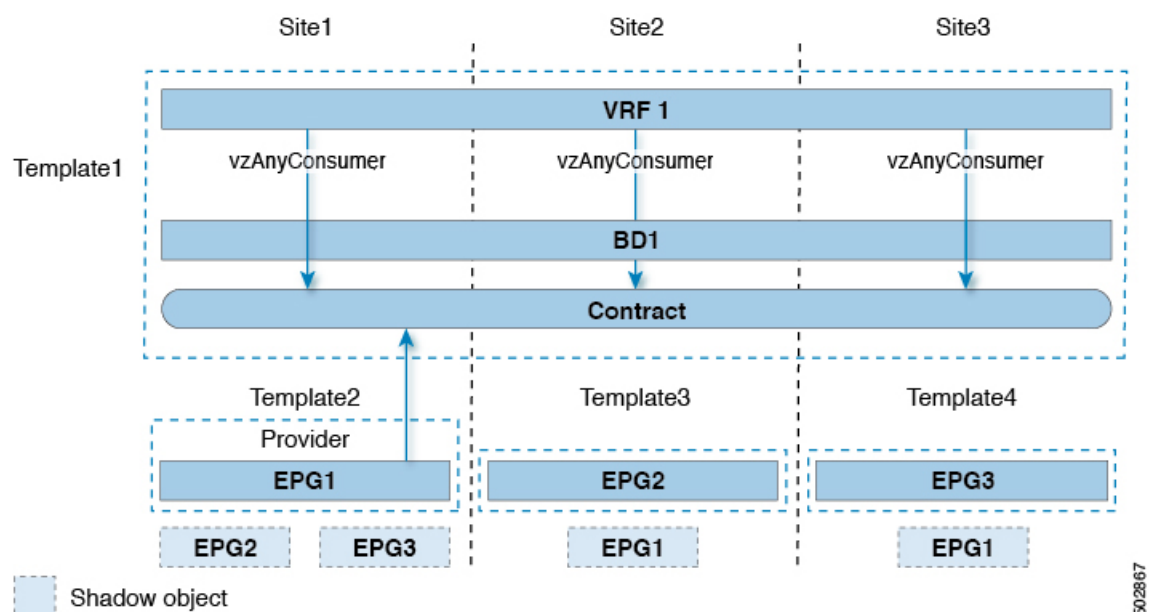
The following example shows intra-VRF communication between a single provider EPG (for example, shared service) and all other EPGs within the same VRF consuming the service.

In this case you will need to create multiple templates:

- A single template for the shared objects (VRF, Contract, BDs) deployed to every site.
- And a separate template for every site combination containing the objects deployed only to those sites.

The following figure shows a single stretched VRF/BD configuration. Alternatively, you can also configure and map a dedicated BD for each EPG, in which case shadow BDs would be deployed in the remote sites.

Figure 42:



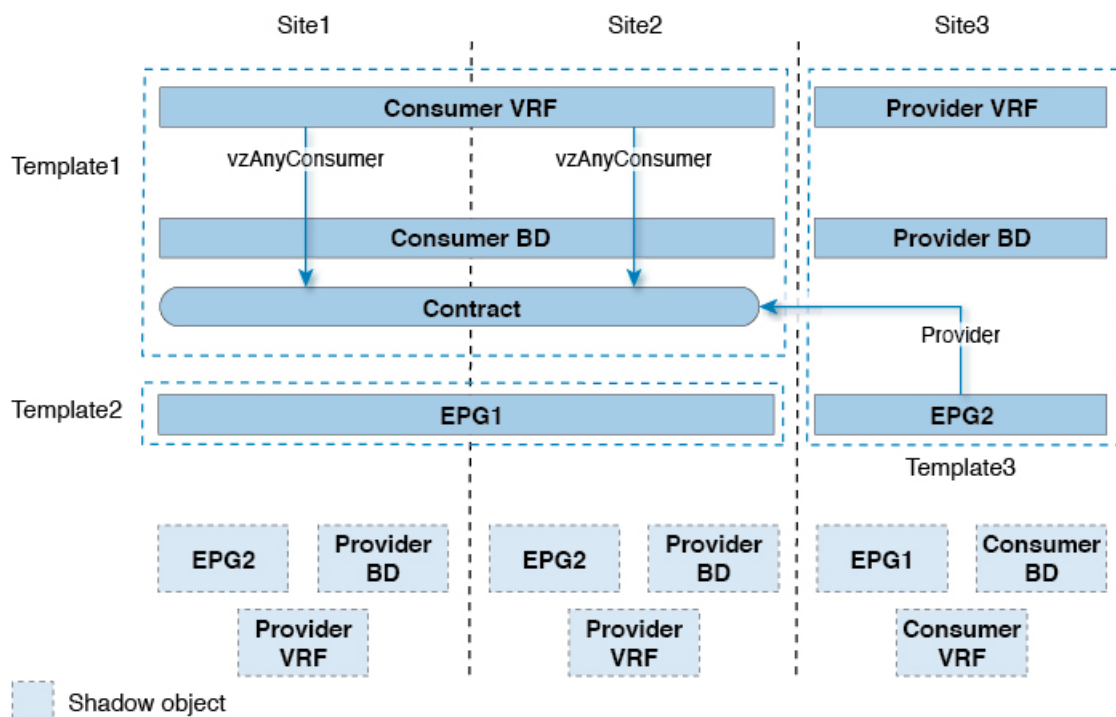
Provider EPG In Its Own VRF

The following example shows communication between a single EPG (for example, shared service provider) in its own VRF and all EPGs within a different, vzAny VRF. The provider EPG can be deployed to the same or a different site as the consumer EPGs in the vzAny VRF.

In this case you will need to create multiple templates:

- A single template for the shared vzAny objects (VRF, Contract, BD) deployed to one or more sites.
- And a separate template for every site combination containing the objects deployed only to those sites.

Figure 43:



502868