



Cisco Nexus Dashboard Deployment and Upgrade Guide, Release 4.3.x

First Published: 2026-08-23

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)

© 2026 Cisco Systems, Inc. All rights reserved.



CONTENTS

Trademarks ?

CHAPTER 1

New and Changed Information 1

New and changed information 1

PART I

Preparing to Deploy Nexus Dashboard 3

CHAPTER 2

Deployment Overview and Requirements 5

Nexus Dashboard deployment overview 5

About supported node types and features 8

CHAPTER 3

Prerequisites and Guidelines 11

General prerequisites and guidelines 11

Prerequisites for the Nexus Dashboard data network and management network 15

Prerequisites for the Nexus Dashboard internal app and service networks 17

Prerequisites for LAN deployments 17

Network prerequisites for LAN deployments 17

Prerequisites for onboarding ACI fabrics in LAN deployments 18

Prerequisites for onboarding NX-OS, IOS XR, and IOS XE devices in LAN deployments 20

Communication ports for LAN deployments 20

Prerequisites for SAN deployments 31

Network prerequisites for SAN deployments 31

Communication ports for SAN deployments 31

Nexus Dashboard persistent IP addresses 40

BGP configuration and persistent IP addresses 46

Round trip time requirements 47

Prerequisites for Nexus Dashboard storage (virtual form factors)	47
Fabric connectivity	49

PART II
Deploying the Cluster 57

CHAPTER 4
Pre-Installation Checklist 59

Pre-installation checklist	59
----------------------------	----

CHAPTER 5
Deploying as a Physical Appliance 63

Prerequisites and guidelines for deploying Nexus Dashboard as a physical appliance	63
Configure a Cisco Integrated Management Controller IP address	64
Enable Serial over LAN in the Cisco Integrated Management Controller	65
Understanding and cabling physical appliances	65
ND-NODE-G5L (UCS-C225-M8)	66
Understanding the ND-NODE-G5L (UCS-C225-M8)	66
Status LEDs and Buttons	69
Physical node cabling	74
ND-NODE-G5S (UCS-C225-M8)	77
Understanding the ND-NODE-G5S (UCS-C225-M8)	77
Physical node cabling	85
ND-NODE-L4 (UCS-C225-M6)	88
Understanding the ND-NODE-L4 (UCS-C225-M6)	88
Physical node cabling	96
SE-NODE-G2 (UCS-C220-M5)	99
Understanding the SE-NODE-G2 (UCS-C220-M5)	99
Physical node cabling	108
Deploy Nexus Dashboard as a physical appliance	110

CHAPTER 6
Deploying in VMware ESX 119

Prerequisites and guidelines for deploying the Nexus Dashboard cluster in VMware ESX	119
Deploy Nexus Dashboard Using VMware vCenter	121
Deploy Nexus Dashboard Directly in VMware ESXi	135

CHAPTER 7
Deploying in Linux KVM 147

Prerequisites and guidelines for deploying the Nexus Dashboard cluster in Linux KVM	147
Verify the I/O latency of a Linux KVM storage device	148
Understanding system resources	149
About the deployment script for QCOW2 images	149
Prerequisites for the deployment script for QCOW2 images	150
Use the deployment script for QCOW2 images	150
Optional post-deployment tasks	156
Deploy Nexus Dashboard in Linux KVM	157

CHAPTER 8

Deploying a Virtual Nexus Dashboard (vND) in Nutanix	169
Prerequisites and guidelines for deploying vNDs in Nutanix	169
Verify the I/O latency of a Nutanix storage device	170
Understanding system resources	171
About the deployment script for QCOW2 images	171
Prerequisites for the deployment script for QCOW2 images	172
Use the deployment script for QCOW2 images	173
Optional post-deployment tasks	179
Install the Nexus Dashboard cluster in Nutanix	180

CHAPTER 9

Deploying a Virtual Nexus Dashboard (vND) in Amazon Web Services (AWS)	203
About hosting a vND on the AWS public cloud	203
Prerequisites and guidelines for deploying the vNDs in Amazon Web Services	205
Prepare Amazon Web Services for the Nexus Dashboard cluster	206
Deploy a virtual Nexus Dashboard (vND) in Amazon Web Services (AWS)	208
Node replacement (RMA)	214

PART III

Upgrading or Migrating to This Release 219

CHAPTER 10

Upgrading a Nexus Dashboard 3.2.2 Cluster to This Release	221
Prerequisites and guidelines for upgrading an existing Nexus Dashboard cluster	221
Nexus Dashboard pre-upgrade validation script	226
(Optional) Convert separate NDFC and NDI clusters to a single co-hosted Nexus Dashboard cluster	227
Supported upgrade paths	228
Upgrade Nexus Dashboard	230

Post-upgrade information and tasks 233

Troubleshooting upgrades 237

CHAPTER 11**Upgrading a Nexus Dashboard 4.1.1 Cluster to This Release 239**

Prerequisites and guidelines for upgrading an existing Nexus Dashboard cluster 239

Nexus Dashboard pre-upgrade validation script 241

Supported upgrade paths 243

Upgrade Nexus Dashboard 245

Troubleshooting upgrades 247

CHAPTER 12**Migrating From DCNM to ND 249**

Migrating from DCNM to ND 249

CHAPTER 13**Migrate from NDB to ND 251**

Migrate from NDB to ND 251



CHAPTER 1

New and Changed Information

- [New and changed information, on page 1](#)

New and changed information

This table provides an overview of the significant changes to the organization and features in this guide from the release in which 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
4.3.1	Support for ESXi 9.0.2	Support is now available for ESXi 9.0.2 when deploying the Nexus Dashboard cluster in VMware ESX: • vND support on vSphere 9.0.2 • VMware integration support for vCenter 9.0.2 For more information, see Deploying in VMware ESX, on page 119 .

Release	New Feature or Update	Where Documented
4.3.1	Nexus Dashboard pre-upgrade validation script as signed plugin process	• Prior to Nexus Dashboard release 4.3.x, the Nexus Dashboard pre-upgrade validation script was a standalone Python script that you would download and run from an external host over SSH to perform health checks across nodes. • Beginning with Nexus Dashboard release 4.3.x, the Nexus Dashboard pre-upgrade validation script is now a signed plugin process, where you download the Cisco-signed .ndp (Nexus Dashboard plugin) file, which Nexus Dashboard validates and runs locally with elevated privileges. For more information, see Nexus Dashboard pre-upgrade validation script, on page 226.
4.3.1	Deployment script for QCOW2 images	Deploy Nexus Dashboard more consistently across KVM-based environments, including RHEL and Nutanix. QCOW2 images now include a deployment validation and orchestration script that simplifies setup, validates compatibility, and promotes consistent deployment across supported platforms For more information, see About the deployment script for QCOW2 images, on page 149.



PART I

Preparing to Deploy Nexus Dashboard

- [Deployment Overview and Requirements, on page 5](#)
- [Prerequisites and Guidelines, on page 11](#)



CHAPTER 2

Deployment Overview and Requirements

- [Nexus Dashboard deployment overview, on page 5](#)

Nexus Dashboard deployment overview

Nexus Dashboard platform

Cisco Nexus Dashboard is a central management console for multiple data center fabrics that provides real-time analytics, visibility, assurance for network policies and operations, as well as policy orchestration for the data center fabrics, such as Cisco ACI and NX-OS.

Nexus Dashboard is the comprehensive management solution for ACI as well as NX-OS deployments spanning LAN fabric, SAN fabric, and IP Fabric for Media (IPFM) networks in data centers powered by Cisco. Nexus Dashboard also supports other devices, such as IOS-XE switches, IOS-XR routers, and non-Cisco devices. Being a multi-fabric controller, Nexus Dashboard manages multiple deployment models such as VXLAN EVPN, classic 3-tier LAN, FabricPath, and routed-based fabrics for LAN while providing ready-to-use control, management, monitoring, and automation capabilities for all these environments. In addition, Nexus Dashboard, when you select SAN installation, Cisco Nexus Dashboard automates Cisco MDS switches and Cisco Nexus-family infrastructure in NX-OS mode with a focus on storage-specific features and analytics capabilities.



Note This document describes how to deploy a Nexus Dashboard cluster initially and onboard the fabrics. After your cluster is up and running, see the Nexus Dashboard [configuration and operation articles](#) for day-to-day operation.

Unified Nexus Dashboard deployment

The Nexus Dashboard (ND) platform and related services were available in the following ways previously:

- Prior to ND release 3.1, Nexus Dashboard shipped with only the platform software and no services included. You would download, install, and enable the services separately (NDI, NDO, and/or NDFC) after the initial ND platform deployment.
- For ND releases 3.1 and 3.2, Nexus Dashboard packaged the ND platform software and the individual services' software in a unified packaging form; however, you still enabled the services separately. Management and Insights of the fabrics were still two independent pieces that were not unified.

In addition, there existed a concept of "deployment mode" in Nexus Dashboard releases 3.1 and 3.2, where you would statically enable specific services in Nexus Dashboard by selecting the deployment mode. However, changing the deployment mode was a disruptive exercise that would wipe out the entire service, including data and reinstalls. And finally, you could not run all the services in a single Nexus Dashboard cluster in Nexus Dashboard releases 3.1 and 3.2.

Beginning with ND release 4.1, the platform and the individual services have been unified into a single product. You no longer deploy and configure the services separately, and you do not have to activate individual services or statically configure deployment modes. In addition, depending on the form factor, Nexus Dashboard allows you to consume any capabilities that were shipped as services in previous releases. The user experience is now unified, as there is no more concept of independent services; instead, all capabilities are now available from a single dashboard view.



Note Depending on the cluster format and the number of cluster nodes that you have deployed, certain features (such as controller, orchestrator, or telemetry) might not be available in the unified Nexus Dashboard product. Review the information in the [Nexus Dashboard Capacity Planning tool](#) to verify what features would be available for your cluster installation.

Hardware vs software stack

Nexus Dashboard is offered as a cluster of specialized Cisco UCS servers (Nexus Dashboard platform) with the software framework (Nexus Dashboard) pre-installed on it. The Cisco Nexus Dashboard software stack can be decoupled from the hardware and deployed in a number of virtual form factors. For the purposes of this document, we will use "Nexus Dashboard hardware" specifically to refer to the hardware and "Nexus Dashboard" to refer to the software stack and the GUI console. In addition, we will use the term "physical Nexus Dashboard" or "pND" to refer to the Cisco UCS physical appliance hardware with the Nexus Dashboard software stack pre-installed on it, and "virtual Nexus Dashboard" or "vND" to refer to the virtual form factor that supports data and app nodes.

This guide describes the initial deployment of the Nexus Dashboard software, which is common for physical and virtual form factors. If you are deploying a physical cluster, see [Nexus Dashboard Hardware Setup Guide](#) for the UCS servers' hardware overview, specification, and racking instructions.



Note Root access to the Nexus Dashboard software is restricted to Cisco TAC only. A special user `rescue-user` is created for all Nexus Dashboard deployments to enable a set of operations and troubleshooting commands. For additional information about the available `rescue-user` commands, see the "Troubleshooting" article in the Nexus Dashboard [documentation library](#).

Available form factors

This release of Cisco Nexus Dashboard can be deployed using a number of different form factors. However, you must use the same form factor for all nodes, mixing nodes of different form factors within the same cluster is not supported. The physical form factor currently supports four different Cisco UCS servers (`SE-NODE-G2`, `ND-NODE-L4`, `ND-NODE-G5S`, and `ND-NODE-G5L`). You can mix `SE-NODE-G2` and `ND-NODE-L4` servers in the same cluster, but you cannot mix a `ND-NODE-G5S` or `ND-NODE-G5L` server in the same cluster as `SE-NODE-G2` and `ND-NODE-L4` servers.

- Physical appliance (.iso) – This form factor refers to the Cisco UCS physical appliance hardware with the Nexus Dashboard software stack pre-installed on it.

The later sections in this document describe how to configure the software stack on the existing physical appliance hardware to deploy the cluster. Setting up the Nexus Dashboard hardware is described in [Nexus Dashboard Hardware Setup Guide](#) for the specific UCS model.

- Virtual Appliance – The virtual form factor allows you to deploy a Nexus Dashboard cluster using VMware ESX (.ova) or RHEL KVM (.qcow2).

The virtual form factor supports the following two profiles:

- Data node – This profile with higher system requirements is designed for higher scale and/or unified deployment.
- App node – This profile with lower system requirements can be deployed as secondary nodes. Can also be deployed as primary nodes but does not support unified deployment.

Support is also available for running a virtual Nexus Dashboard (vND) on the AWS public cloud. See [Deploying a Virtual Nexus Dashboard \(vND\) in Amazon Web Services \(AWS\), on page 203](#) for more information.

Beginning with Nexus Dashboard release 4.2(1), support is also available for running a virtual Nexus Dashboard (vND) on the Nutanix. See [Deploying a Virtual Nexus Dashboard \(vND\) in Nutanix, on page 169](#) for more information.



Note When planning your deployment, ensure to check the list of "Prerequisites and Guidelines" in one of the following sections of this document specific to the form factor you are deploying. A quick reference of the supported form factors, scale, and cluster sizing requirements are available in the [Nexus Dashboard Cluster Sizing](#) tool.

Scale and cluster sizing guidelines

A basic Nexus Dashboard deployment typically consists of 1 or 3 `primary` nodes, which are required for the cluster to come up. Depending on scale requirements, 3-node or larger clusters can be extended with up to 3 additional `secondary` nodes to support higher scale. Adding secondary nodes to the cluster must be done during cluster deployment and cannot be added once the cluster has been initialized.

- For physical clusters, you can also add up to 2 `standby` nodes for easy cluster recovery in case of a primary node failure.
- For virtual clusters, up to 2 `standby` nodes are also supported, but only with a 3-node vND (app) profile for a Controller-only or Orchestration-only deployment.

Exact number of additional secondary nodes required for your specific use case is available from the [Nexus Dashboard Cluster Sizing](#) tool.

Scale and cluster sizing limitations

These limitations apply to scale and cluster sizing:

- The `ND-NODE-G5L` (UCS-C225-M8) node is available only in a 3-node cluster configuration.

- Deploying a 4-cluster node, where the cluster consists of:
 - 3 virtual nodes (data), and
 - 1 standby node

is not a supported configuration. Redeploy this cluster without the standby node. If one of the nodes in the cluster fails, reinstall a new node and navigate to **Admin > System Status > Nodes**, then click **Actions > Re-Register** to add the reinstalled node back into the cluster.

- Single-node deployments cannot be extended to a 3-node cluster after the initial deployment.
If you deploy a single-node cluster and want to extend it to a 3-node cluster or add `secondary` nodes, you will need to back it up, deploy a new 3-node base cluster, and then restore the backup on this later. For more information, see [Backing Up and Restoring Your Nexus Dashboard](#)
- Single-node deployments do not support additional `secondary` or `standby` nodes.
- For 3-node clusters, at least two `primary` nodes are required for the cluster to remain operational.
For more information, see [Deploying Highly Available Services with Cisco Nexus Dashboard](#).
- You cannot dynamically add or remove primary or secondary nodes to change the overall size of an existing cluster. To resize a cluster, you must back up the cluster, rebuild it to the desired size, and then restore the configuration. Ensure that the new cluster size meets the supported scale and feature requirements for your deployment.
- When a cluster size change is required, delete only one secondary node at a time. Wait until the cluster returns to a healthy state before you delete the next secondary node. Verify the cluster health in the Nexus Dashboard GUI or by running the `acs health` command. Do not delete multiple secondary nodes concurrently or while the cluster is converging.
- For validated scale, all required worker nodes should be included when the Nexus Dashboard cluster is initially deployed. Additional worker nodes can also be added to an existing cluster; however, in the current release, adding worker nodes after initial deployment does not dynamically increase the cluster's validated scale capacity. The additional nodes can still help distribute workloads across a larger set of resources, reducing concentration on individual nodes. If increased scale capacity is required after worker nodes have been added, the cluster can be backed up and restored so that the expanded topology is incorporated into the scale configuration.

About supported node types and features

- Several node types are available for Nexus Dashboard, including SE-NODE-G2 (UCS-C220-M5), ND-NODE-L4 (UCS-C225-M6), ND-NODE-G5S (UCS-C225-M8), and ND-NODE-G5L (UCS-C225-M8).
- Key features that can be enabled per fabric include **Controller**, **Telemetry**, and **Orchestration**.
- Node and feature compatibility is subject to guidelines and limitations, such as restrictions on mixing node types and enabling all features on certain clusters.

Supported Node Types and Feature Overview

These node types are available for Nexus Dashboard:

- `ND-NODE-G5L` (UCS-C225-M8). The product ID of the 3-node cluster is `ND-CLUSTERG5L`.
- `ND-NODE-G5S` (UCS-C225-M8). The product ID of the 3-node cluster is `ND-CLUSTERG5S`.
- `ND-NODE-L4` (UCS-C225-M6). The product ID of the 3-node cluster is `ND-CLUSTER-L4`.
- `SE-NODE-G2` (UCS-C220-M5). The product ID of the 3-node cluster is `SE-CL-L3`.



Note Support for the `SE-NODE-G2` (UCS-C220-M5) node will come to an end in an upcoming Nexus Dashboard release. The EOL milestones for this appliance can be found here:

[End-of-Sale and End-of-Life Announcement for the Cisco UCS M5 Rack Servers \(C220 M5, C240 M5, C480 M5\).](#)

Future Nexus Dashboard release notes will contain details about the last release where this appliance will be supported.

In addition, in LAN deployments, these are the available features that you can leverage:

- **Controller** : Also referred to as Fabric Management. This feature is used to manage NX-OS and non-NX-OS switches (such as Catalyst, ASR, and so on). This includes creating any non-ACI fabric types, as well as performing software upgrades and creating new configurations on those fabrics.
- **Telemetry** : This feature provides telemetry functionality, similar to the functionality provided by Nexus Dashboard Insights in prior releases. You can enable and use the **Telemetry** feature when you create or edit a fabric through **Manage > Fabrics** .
- **Orchestration** : You can use the **Orchestration** feature through Nexus Dashboard to connect multiple ACI fabrics together, and consolidate and deploy tenants, along with network and policy configurations, across multiple ACI fabrics. You can enable and use the **Orchestration** feature when you add an ACI through **Admin > System Settings > Multi-cluster connectivity > Connect Cluster**.

For each fabric managed by Nexus Dashboard, you can enable the following features independently within the same Nexus Dashboard cluster or across a multi-cluster federated Nexus Dashboard deployment:

- Controller
- Telemetry
- Orchestration



Note Enabling all capabilities on a single cluster might not be available in some cluster deployment formats. For a quick reference of the supported form factors, scale, and cluster sizing requirements specific to your Nexus Dashboard deployment, see the [Nexus Dashboard Capacity Planning tool](#).

Guidelines and limitations:

- You cannot mix the newer `ND-NODE-G5L` or `ND-NODE-G5S` (UCS-C225-M8) nodes in a cluster with the older `SE-NODE-G2` (UCS-C220-M5) and `ND-NODE-L4` (UCS-C225-M6) nodes.

In addition, you can only have homogeneous `ND-NODE-G5L` or `ND-NODE-G5S` clusters. In other words:

- A cluster containing a ND-NODE-G5L node can only have additional ND-NODE-G5L nodes in that cluster, and
- A cluster containing a ND-NODE-G5S node can only have additional ND-NODE-G5S nodes in that cluster.
- The ND-NODE-G5L node is only supported for LAN deployments.
- The virtual form factor does not support all features in many cluster sizes and types, as described in the [Cisco Nexus Dashboard Verified Scalability Guide](#).
- If a Nexus Dashboard node becomes unavailable, some Telemetry functionality may be impacted. Telemetry dynamic jobs such as Assurance, Bug Scan, and Log Collection, are blocked and will not run until the node is restored.
- Telemetry functionality is fully supported during a node failure when the deployment is operating at 50% or less of the supported scale for its form factor. For example, a 3vND data form factor supports up to 100 switches at full scale. If a deployment contains 50 switches and a Nexus Dashboard node becomes unavailable, Telemetry services will continue to operate normally. For deployments operating above 50% of the supported scale, Telemetry behavior during a node failure is provided on a best-effort basis, and some functionality may be degraded.

Windows 11 VDI client sizing

In Windows 11 virtual desktop infrastructure (VDI) environments, allocate sufficient client resources to provide stable Nexus Dashboard UI performance when multiple users access the UI concurrently. The following configurations have been tested and validated for each concurrent VDI user session.

These sizing requirements apply to VDI client sessions used to access the Nexus Dashboard UI.

Component	Minimum supported configuration per user	Recommended configuration per user
CPU	8 vCPUs	16 vCPUs
Processor speed	2.1 GHz or higher	2.1 GHz or higher
Processor	Compatible 64-bit processor with two or more cores; Intel Xeon Silver 4116 or equivalent	Compatible 64-bit processor with two or more cores; Intel Xeon Silver 4116 or equivalent
Memory	8 GB RAM	32 GB RAM
Hardware acceleration	Enabled or disabled	Enabled
Display resolution	Higher than 1280 × 720 pixels	Higher than 1280 × 720 pixels



CHAPTER 3

Prerequisites and Guidelines

- [General prerequisites and guidelines, on page 11](#)
- [Prerequisites for the Nexus Dashboard data network and management network, on page 15](#)
- [Prerequisites for the Nexus Dashboard internal app and service networks, on page 17](#)
- [Prerequisites for LAN deployments, on page 17](#)
- [Prerequisites for SAN deployments, on page 31](#)
- [Nexus Dashboard persistent IP addresses, on page 40](#)
- [Round trip time requirements, on page 47](#)
- [Prerequisites for Nexus Dashboard storage \(virtual form factors\), on page 47](#)
- [Fabric connectivity, on page 49](#)

General prerequisites and guidelines

This section describes requirements and guidelines for the Nexus Dashboard cluster regardless of the deployment type.

General deployment guidelines and restrictions

- Deploying virtual Nexus Dashboard VMs on remote storage is unsupported and may lead to unexpected behavior.

Domain Name System (DNS) and Network Time Protocol (NTP)

The Nexus Dashboard nodes require valid DNS and NTP servers for all deployments and upgrades.

Lack of valid DNS connectivity (such as if using an unreachable or a placeholder IP address) can prevent the system from deploying or upgrading successfully, as well as impact regular services functionality.



Note Nexus Dashboard acts as both a DNS client and resolver. It uses an internal Core DNS server which acts as DNS resolver for internal services. It also acts as a DNS client to reach external hosts within the intranet or the Internet, hence it requires an external DNS server to be configured.

The following guidelines apply for DNS:

- Prior to Nexus Dashboard release 4.3.1, the DNS domain name was created using the cluster name. Starting with Nexus Dashboard release 4.3.1, the DNS name is fixed to `cluster.case.local`.
- For external DNS servers, both TCP and UDP traffic must be allowed. See [Communication ports for LAN deployments, on page 20](#) and [Communication ports for SAN deployments, on page 31](#) for more information.
- Nexus Dashboard does not support DNS servers with wildcard records.

Nexus Dashboard also supports NTP authentication using symmetrical keys. If you want to enable NTP authentication, you will need to provide the following information during cluster configuration:

- **NTP Key**—A cryptographic key that is used to authenticate the NTP traffic between the Nexus Dashboard and the NTP server(s). You will define the NTP servers in the following step, and multiple NTP servers can use the same NTP key.
- **Key ID**—Each NTP key must be assigned a unique key ID, which is used to identify the appropriate key to use when verifying the NTP packet.
- **Auth Type**—This release supports MD5, SHA, and AES128CMAC authentication types.

The following guidelines apply when enabling NTP authentication:

- We recommend that you do not use a Windows server as the NTP server.
- For symmetrical authentication, any key you want to use must be configured the same on both your NTP server and Nexus Dashboard.
The ID, authentication type, and the key/passphrase itself must match and be trusted on both your NTP server and Nexus Dashboard.
- Multiple servers can use the same key.
In this case the key must only be configured once on Nexus Dashboard, then assigned to multiple servers.
- Both Nexus Dashboard and the NTP servers can have multiple keys as long as key IDs are unique.
- This release supports SHA1, MD5, and AES128CMAC authentication/encoding types for NTP keys.



Note We recommend using AES128CMAC due to its higher security.

- When adding NTP keys in Nexus Dashboard, you must tag them as `trusted`; untrusted keys will fail authentication.
This option allows you to easily disable a specific key in Nexus Dashboard if the key becomes compromised.
- You can choose to tag some NTP servers as `preferred` in Nexus Dashboard.
NTP clients can estimate the "quality" of an NTP server over time by taking into account RTT, time response variance, and other variables. Preferred servers will have higher priority when choosing a primary server.
- If you are using an NTP server running `ntpd`, we recommend version 4.2.8p12 at a minimum.
- The following restrictions apply to all NTP keys:

- The maximum key length for SHA1 and MD5 keys is 40 characters, while the maximum length for AES128 keys is 32 characters.
- Keys that are shorter than 20 characters can contain any ASCII character excluding '#' and spaces. Keys that are over 20 characters in length must be in hexadecimal format.
- Keys IDs must be in the 1-65535 range.
- If you configure keys for any one NTP server, you must also configure the keys for all other servers.
- Nexus Dashboard nodes must be in synchronization with the NTP server; however, there can be latency of up to 1 second between the Nexus Dashboard nodes. If the latency is greater than or equal to 1 second between the Nexus Dashboard nodes, this may result in unreliable operations on the Nexus Dashboard cluster.
- These are the requirements for NTP delay, offset, and jitter:
 - Delay: < 100 ms
 - Offset: ±25 ms
 - Jitter: < 10 ms

Enabling and configuring NTP authentication is described as part of the deployment steps in the later sections.

IPv4 and IPv6 support

Nexus Dashboard supports IPv4-only, IPv6-only or dual-stack configurations for the cluster nodes and services.

When defining an IP address configuration, the following guidelines apply:

- All nodes and networks in the cluster must have a uniform IP configuration, either IPv4-only, IPv6-only, or dual stack IPv4/IPv6.
- If you deploy the cluster in IPv4-only mode and want to switch to dual stack IPv4/IPv6 or IPv6-only, you must redeploy the cluster.
- For dual stack configurations:
 - The data, management, app, and service networks must be in dual stack mode.
Mixed configurations, such as IPv4 data network and dual stack management network, are not supported.
 - For IPv6-based Nexus Dashboard deployments, the CIMCs of all physical servers must also have IPv6 addresses.
 - You can configure either IPv4 or IPv6 addresses for the nodes' management network during initial node bring up, but you must provide both types of IP addresses during the cluster bootstrap workflow.
Management IP addresses are used to log in to the nodes for the first time to initiate cluster bootstrap process.
 - Kubernetes internal core services will start in IPv4 mode.
 - DNS will serve and forward both IPv4 and IPv6 requests.
 - VXLAN overlay for peer connectivity will use data network's IPv4 addresses.

Both IPv4 and IPv6 packets are encapsulated within the VXLAN's IPv4 packets.

- The GUI will be accessible on both IPv4 and IPv6 management network addresses, provided both are configured.
- For IPv6-only configurations:
 - IPv6-only mode is supported for physical and virtual form factors only.
Clusters deployed through the vND deployment process on AWS public cloud do not support IPv6-only or dual stack mode.
 - You must provide IPv6 management network addresses when initially configuring the nodes.
After the nodes are up, these IP addresses are used to log in to the GUI and continue cluster bootstrap process.
 - You must provide IPv6 CIDRs for the internal app and service networks described above.
 - You must provide IPv6 addresses and gateways for the data and management networks described above.
 - All internal services will start in IPv6 mode.
 - VXLAN overlay for peer connectivity will use data network's IPv6 addresses.
IPv6 packets are encapsulated within the VXLAN's IPv6 packets.
 - All internal services will use IPv6 addresses.
 - IPv6 addresses are required for physical servers' CIMCs.

Necessary URLs for certain connections

There are certain URLs that Nexus Dashboard must reach that are necessary for these connections:

- Cisco Intersight: Connecting your Nexus Dashboard cluster to Cisco Intersight has these benefits:
 - Automatic meta data updates that certain features can use to provide updated data
 - TAC log collection and uploads
- Connecting to Smart Licensing
- Pulling energy management stats from electricity maps

These are the URLs that Nexus Dashboard must reach for these connections and why:

URL	Protocol/Port/Service	Description
amazontrust.com	TCP/80(HTTP) TCP/443(HTTPS)	Used to securely connect to Cisco Intersight
connectdna.cisco.com	TCP/443(HTTPS)	Used to securely connect to Cisco Intersight and Smart Licensing
swapi.cisco.com	TCP/443(HTTPS)	Used to securely connect to Cisco Smart Licensing
svc.ucs-connect.com	TCP/443(HTTPS)	Used to securely connect to Cisco Intersight

URL	Protocol/Port/Service	Description
svc-static1.ucs-connect.com	TCP/443(HTTPS)	Used to securely connect to Cisco Intersight
svc.eu-central-1.intersight.com	TCP/443(HTTPS)	Used to securely connect to Cisco Intersight (EMEA Region)
svc-static1.eu-central-1.intersight.com	TCP/443(HTTPS)	Used to securely connect to Cisco Intersight (EMEA Region)

Prerequisites for the Nexus Dashboard data network and management network

Nexus Dashboard is deployed as a cluster, connecting each node to two networks. When first configuring Nexus Dashboard, for each cluster node, you will need to provide two IP addresses for the two Nexus Dashboard interfaces:

- One connected to the data network, which is used for back-end, cluster, and Infra connectivity for optimal performance
- The other connected to the management network, which is used for seamless GUI and front-end operations

Table 2: External network purpose

Data network	Management network
• Nexus Dashboard node clustering • Service to service communication • Nexus Dashboard nodes to Cisco APIC and NX-OS controller capability communication • Telemetry traffic for switches and on-boarded fabrics	• Accessing Nexus Dashboard GUI • Accessing Nexus Dashboard CLI using SSH • DNS and NTP communication • Nexus Dashboard firmware upload • Intersight device connector • AAA traffic • Multi-cluster connectivity

The two networks have the following requirements:

- The management network and data network must be in different subnets.



Note Nexus Dashboard management interface (`bond1`) has internal iptables rules that rate-limit ICMP packets to an average of 6 packets per second with a burst limit of 5. This ICMP rate and burst limit also applies to the data network port (`bond0`). If you are using ICMP-based monitoring tools to track the health of the management network, you may observe intermittent packet drops if the polling frequency exceeds these limits. This is expected behavior designed to protect the management plane.

- Changing the data subnet requires re-deploying the cluster, so we recommend using a larger subnet (such as /27) than the bare minimum required by the nodes and features to account for any additional features that may require more IP addresses in the future.
- When setting up remote authentication, AAA server must not be in the same subnet as the data interface.
- For physical clusters, the management network must provide IP reachability to each node's CIMC using TCP ports 22 and 443 as the Nexus Dashboard cluster configuration uses each node's CIMC IP address to configure the node.
- The data network interface requires a minimum MTU of 1500 to be available for the Nexus Dashboard traffic.

Higher MTU can be configured if desired on the switches to which the nodes are connected.



Note If external VLAN tag is configured for switch ports that are used for data network traffic, you must enable jumbo frames or configure custom MTU equal to or greater than 1504 bytes on the switch ports where the nodes are connected.

- If you are using telemetry, by default, the data network must provide IP reachability to the in-band network of each fabric and of the Cisco APIC for an ACI fabric (if you are using the orchestration functionality), as well as to these integrations.



Note You can also define routes in the route table of the Nexus Dashboard and use the management network instead to reach to any of the following services.

- For DNS integration, to the DNS server.
- For Panduit PDU integration, to the Panduit PDU server.
- For External Kafka integration, to the External Kafka server (consumer).
- For SysLog integration, to the SysLog server.
- For Network-Attached Storage integration, to the Network-Attached Storage server.
- For VMware vCenter integration, to the VMware vCenter.
- For AppDynamics integration, to the AppDynamics controller.

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



Note If all the integrations are in same subnet as the management network, then they will use the management network.

Prerequisites for the Nexus Dashboard internal app and service networks

Two additional internal networks are required for communication between the containers used by the Nexus Dashboard:

- App network--Used for applications internally within Nexus Dashboard. The app network must be a /16 network for IPv4 or /108 network for IPv6 and a default value is pre-populated during deployment.
- Service network--Used internally by the Nexus Dashboard. The service network must be a /16 network for IPv4 or /108 network for IPv6 and a default value is pre-populated during deployment.

If you are planning to deploy multiple Nexus Dashboard clusters, they can use the same application and service subnets.

**Note**

Communications between containers deployed in different Nexus Dashboard nodes is VXLAN-encapsulated and uses the data interfaces IP addresses as source and destination. This means that the app network and service network addresses are never exposed outside the data network and any traffic on these subnets is routed internally and does not leave the cluster nodes.

For example, if you had another service (such as DNS) on the same subnet as the app or service network, you would not be able to access it from your Nexus Dashboard as the traffic on that subnet would never be routed outside the cluster. As such, when configuring these networks, ensure that they are unique and do not overlap with any existing networks or services external to the cluster, which you may need to access from the Nexus Dashboard cluster nodes.

For the same reason, we recommend not using 169.254.0.0/16 (the Kubernetes `br1` subnet) for the app or service subnets.

Prerequisites for LAN deployments

Network prerequisites for LAN deployments

These network prerequisites apply for LAN deployments:

- All new Nexus Dashboard deployments must have the management network and data network in different subnets.
- Interfaces on both data and management networks can be either Layer 2 or Layer 3 adjacent. For data network Layer 3 adjacency, you must configure BGP during the bootstrap process. Management network interfaces do not support the BGP protocol. If different Nexus Dashboard nodes are deployed with management addresses in different subnets, those will simply be routed to one other.
- You must use persistent data IP addresses to bring up the cluster, so you must allocate a certain number of persistent IP addresses depending on your configuration.
 - If your cluster has 1 node, allocate 3 persistent IP addresses.

- If your cluster has 3 or more nodes, allocate 5 persistent IP addresses.
- If you configure dual stack IPv4 and IPv6, then add the same number of persistent IP addresses for IPv6 (in other words, 5 IPv4 and 5 IPv6 persistent IP addresses if you configure dual stack).

For more information about persistent IP addresses, see [Nexus Dashboard persistent IP addresses, on page 40](#). You must allocate the minimum required persistent IP addresses during the bootstrap process. You can allocate additional persistent IP addresses after the cluster is deployed using the External Service Pools configuration in the GUI.

Prerequisites for onboarding ACI fabrics in LAN deployments

These network prerequisites apply for onboarding ACI fabrics in LAN deployments:

- If you are planning to use orchestration to manage Cisco ACI fabrics, you can establish connectivity from either the data interface or the management interface to either the in-band or out-of-band (OOB) interface of each fabric's APIC cluster or both.

If the fabric connectivity is from the Nexus Dashboard's management interface, you must configure specific static routes or ensure that the management interface is part of the same IP subnet of the APIC interfaces.

Additional prerequisites for using orchestration with ACI fabrics

If you plan to use orchestration with ACI fabrics, these prerequisites also apply:

- If you plan to use orchestration with ACI fabrics and remote leaf switches, these restrictions apply:
 - Remote leaf switches in one fabric cannot use another fabric's L3Out.
 - Stretching a bridge domain between one fabric (local leaf or remote leaf) and a remote leaf in another fabric is not supported.
- Orchestration is only supported on single-node Nexus Dashboard clusters (virtual-data profile or physical appliances) for non-production (lab) deployments. If you want to enable Orchestration on one of these form factors, it must be enabled using the built-in swagger API.
 1. From the Nexus Dashboard UI, click on the "?" icon and choose **Help Center**.
 2. In the **Help Center**, click on **API reference: Swagger (In-product)**.
 3. Within the API listing, click the **Infra** group from the left navigation.
 4. Locate the **System Settings** sub-menu and click the arrow to expand it, if necessary, then search for `/settings/general/actions/enableOrchestration`.
 5. Expand the API and click **Try it Out**.

Orchestration services will now be enabled on your cluster.

Additional prerequisites for using telemetry with ACI fabrics

If you plan to use telemetry with ACI fabrics, these prerequisites also apply:

- Telemetry collection is supported with OOB networks of APIC and switches as long as they're running ACI version 6.1(2f) or later.
- You have configured NTP settings on Cisco APIC.

For more information, see [Configure NTP in ACI Fabric Solution](#).

- If you plan to use flow telemetry, the **Node Control** policy on ACI must be set to **Telemetry Priority**. If you are using traffic analytics, then you can use Netflow or Telemetry priority:
 - Using Netflow priority mode allows ACI to also export standard netflow records from ACI to an external source, such as Stealthwatch.
 - Telemetry priority supports exporting netflow records for Traffic Analytics to Nexus Dashboard only.

- If you plan to use the flow telemetry or traffic analytics functions, Precision Time Protocol (PTP) must be enabled on Cisco APIC so that telemetry can correlate flows from multiple switches accordingly

In Cisco APIC, choose **System > System Settings > PTP and Latency Measurement > Admin State** to enable PTP.

The quality of the time synchronization using PTP depends on the accuracy of the PTP Grandmaster (GM) clock which is the source of the clock, and the accuracy and the number of PTP devices such as ACI switches and IPN devices in between.

Although a PTP GM device is generally equipped with a GNSS/GPS source to achieve the nanosecond accuracy which is the standard requirement of PTP, microsecond accuracy is sufficient for flow telemetry, hence a GNSS/GPS source is typically not required.

For a single-pod ACI fabric, you can connect your PTP GM using leaf switches. Otherwise, one of the spine switches will be elected as a GM. For a multi-pod ACI fabric, you can connect your PTP GM using leaf switches or using IPN devices. Your IPN devices should be PTP boundary clocks or PTP transparent clocks so that ACI switch nodes can synchronize their clock across pods. To maintain the same degree of accuracy across pods, it is recommended to connect your PTP GM using IPN devices.

See section "Precision Time Protocol" in the *Cisco APIC System Management Configuration Guide* for details about PTP connectivity options.

- For inband telemetry, you have configured in-band management as described in [Cisco APIC and Static Management Access](#). From ACI release 6.1(2) and later, you can also use out-of-band telemetry. Flow telemetry is supported only with in-band management but traffic analytics is supported on either.
- If one or more DNS Domains are set under DNS Profiles, it is mandatory to set one DNS Domain as default.

In Cisco APIC, choose **Fabric > Fabric Policies > Policies > Global > DNS Profile > default > DNS Domains** and set one as default.

Failure to do so will result in the same switch appearing multiple times in the telemetry Flow map.

- Deploy ACI in-band network by configuring EPG using the following:
 - Tenant = `mgmt`
 - VRF = `inb`
 - BD = `inb`
 - Node Management EPG = `default/<any_epg_name>`

- Nexus Dashboard's data-network IP address and ACI fabric's in-band IP address must be in different subnets.

Prerequisites for onboarding NX-OS, IOS XR, and IOS XE devices in LAN deployments

See [Deploying Nexus Fabrics with Telemetry on Cisco Nexus Dashboard](#) for additional useful information.

These network prerequisites apply for onboarding NX-OS, IOS XR, and IOS XE devices in LAN deployments:

- If you are planning to use orchestration to manage NX-OS fabrics, the data network must have in-band reachability for NX-OS fabrics.

Additional prerequisites for using telemetry with NX-OS fabrics or standalone NX-OS switches

If you plan to use telemetry with NX-OS fabrics or standalone NX-OS switches, these prerequisites also apply:

- The data network must have IP reachability to the fabrics' in-band or out-of-band IP addresses.



Note If you are using the Flow Telemetry feature, the data network must have IP reachability to the fabric's in-band IP addresses.

- To enable Flow Telemetry or Traffic Analytics, Precision Time Protocol (PTP) must be configured on all nodes you want to support with telemetry.

In both managed and monitor fabric mode, you must ensure PTP is correctly configured on all nodes in the fabric. You can enable PTP in the fabric setup's **Advanced** tab by checking the **Enable Precision Time Protocol (PTP)** option.

The PTP grandmaster clock should be provided by a device that is external to the network fabric.

The quality of the time synchronization using PTP depends on the accuracy of the PTP Grandmaster (GM) clock which is the source of the clock, and the accuracy and the number of PTP devices along the network path. Although a PTP GM device is generally equipped with a GNSS/GPS source to achieve the nanosecond accuracy which is the standard requirement of PTP, microsecond accuracy is sufficient for flow telemetry, hence a GNSS/GPS source is typically not required.

For details about manually configuring Precision Time Protocol on Nexus switches, see [Cisco Nexus 9000 Series NX-OS System Management Configuration Guide](#).

Communication ports for LAN deployments

Nexus Dashboard uses TLS or mTLS with encryption to protect data privacy and integrity while in transit.

Open all ports described in the tables in this section.

This table lists the management network communication ports for LAN deployments, where in the **Direction** column:

- **In** means toward the cluster

- **Out** means from the cluster toward the fabric or outside world

Table 3: Management network communication ports for LAN deployments

Service	Port	Protocol	Direction (In /Out)	Connection
Nexus Dashboard agent port	30022	TCP	In/Out	When managing SONiC devices supported by Nexus Dashboard, this port must be open towards both management and data networks of the Nexus Dashboard. Nexus Dashboard installs an agent on the SONiC device that opens a web-socket to the 30022 port.
ICMP	ICMP	ICMP	In/Out	Other cluster nodes, CIMC, default gateway, switch discovery. ICMP is required for the cluster to be operational so you must allow the ICMP service. Note Adding or discovering LAN devices uses ICMP echo packets as part of the discovery process. So if you have a firewall between the Nexus Dashboard cluster and your switches, it must allow ICMP messages through or the discovery process will fail. ICMP traffic on the management interface is rate-limited to an average of 6 packets/sec and a burst of 5. Monitoring systems should be configured with this limit in mind to avoid false-positive alerts regarding packet loss.
DHCP	67	UDP	In	If the local DHCP server is configured for bootstrap or POAP purposes. Note When using Nexus Dashboard as a local DHCP server for POAP purposes, all Nexus Dashboard primary node IP addresses must be configured as DHCP relays. Whether the Nexus Dashboard nodes' management IP addresses are bound to the DHCP server is determined by the LAN Device Management Connectivity in the server settings.
DHCP	68	UDP	Out	
DNS	53	TCP/UDP	Out	DNS server
HTTP	80	TCP	Out	Internet/proxy

Service	Port	Protocol	Direction (In /Out)	Connection
HTTP (PnP)	9666	TCP	In	Cisco Plug and Play (PnP) for Catalyst devices is accomplished using Nexus Dashboard HTTP port 9666 and HTTPS port 9667. HTTP on port 9666 is used to send CA certificate bundle to devices to prime the device for HTTPS mode and actual PnP happens over HTTPS on port 9667 afterwards. PnP service, as with POAP, runs on a persistent IP address that is associated with either the management or data subnet. The persistent IP subnet is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings.
HTTP (POAP)	80	TCP	In	Only used for device zero-touch provisioning using POAP, where devices can send (limited jailed write-only access to Nexus Dashboard) basic inventory information to Nexus Dashboard to start secure POAP communication. Nexus Dashboard Bootstrap or POAP can be configured for TFTP or HTTP/HTTPS. The SCP-POAP service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings.
HTTPS	443	TCP	In/Out	UI, other clusters (for multi-cluster connectivity), fabrics, Internet/proxy
HTTPS/HTTP (NX-API)	443/80	TCP	Out	NX-API HTTPS/HTTP client connects to device NX-API server on port 443/80, which is also configurable. NX-API is an optional feature, used by limited set of Nexus Dashboard functions.

Service	Port	Protocol	Direction (In/Out)	Connection
HTTPS (PnP)	9667	TCP	In	Cisco Plug and Play (PnP) for Catalyst devices is accomplished using Nexus Dashboard HTTP port 9666 and HTTPS port 9667. HTTP on port 9666 is used to send CA certificate bundle to devices to prime the device for HTTPS mode and actual PnP happens over HTTPS on port 9667 afterwards. PnP service, as with POAP, runs on a persistent IP address that is associated with either the management or data subnet. The persistent IP subnet is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings.
HTTPS (POAP)	443	TCP	In	Secure POAP is accomplished using the Nexus Dashboard HTTPS Server on port 443. The HTTPS server is bound to the SCP-POAP service and uses the same persistent IP address assigned to that pod. The SCP-POAP service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings.
Infra-Service	30012 30021 30500-30600	TCP/UDP	In/Out	Other cluster nodes
KMS	9880	TCP	In/Out	Other cluster nodes and ACI fabrics
LDAP	389 636	TCP	Out	LDAP server
NTP	123	UDP	Out	NTP server
NX-API	8443	TCP	In/Out	Used by Cisco MDS 9000 Series switches with NX-OS release 9.x and later for performance monitoring.
Radius	1812	TCP	Out	Radius server

Service	Port	Protocol	Direction (In /Out)	Connection
SCP	22	TCP	In/Out	SCP is used by various features to transfer files between devices and Nexus Dashboard, such as for archiving backup files to remote server.. The Nexus Dashboard SCP service serves as the SCP server for both downloads and uploads. SCP is also used by the POAP client on the devices to download POAP-related files. The SCP-POAP service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings.
SCP/Show Techcollection	22	TCP	Out	Transport tech-support file from persistent IP address of Nexus Dashboard POAP-SCP pod to a separate ND cluster running telemetry. The SCP-POAP service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings
SMTP	25	TCP	Out	You can configure the SMTP port on the Admin > Server Settings > General page. This is an optional feature.
SNMP	161	TCP/UDP	Out	SNMP traffic from Nexus Dashboard to devices.
SNMP Trap	2162	UDP	In	SNMP traps from devices to Nexus Dashboard are sent out toward the persistent IP address associated with the SNMP-Trap/Syslog service pod. The SNMP-Trap-Syslog service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings,
SSH	22	TCP	In/Out	CLI and CIMC of the cluster nodes

Service	Port	Protocol	Direction (In/Out)	Connection
TAC Assist	8884	TCP	In/Out	Other cluster nodes Used for TAC Assist, which is a service to collect show tech from switches and upload the information to Intersight. This port is used to exchange show tech data across cluster nodes.
TACACS	49	TCP	Out	TACACS server
TFTP (POAP)	69	UDP	In	Only used for device zero-touch provisioning using POAP, where devices can send (limited jailed write-only access to Nexus Dashboard) basic inventory information to Nexus Dashboard to start secure POAP communication. Nexus Dashboard bootstrap or POAP can be configured for TFTP or HTTP/HTTPS. The SCP-POAP service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings.

This table lists the data network communication ports for LAN deployments, where in the **Direction** column:

- **In** means toward the cluster
- **Out** means from the cluster toward the fabric or outside world

Table 4: Data network communication ports for LAN deployments

Service	Port	Protocol	Direction (In/Out)	Connection
Nexus Dashboard agent port	30022	TCP	In/Out	When managing SONiC devices supported by Nexus Dashboard, this port must be open towards both management and data networks of the Nexus Dashboard. Nexus Dashboard installs an agent on the SONiC device that opens a web-socket to the 30022 port.

Service	Port	Protocol	Direction (In /Out)	Connection
DHCP	67	UDP	In	If the Nexus Dashboard local DHCP server is configured for bootstrap or POAP purposes. Note When using Nexus Dashboard as a local DHCP server for POAP purposes, all Nexus Dashboard primary node IP addresses must be configured as DHCP relays. Whether the Nexus Dashboard nodes' data IP addresses are bound to the DHCP server is determined by the LAN Device Management Connectivity in the server settings.
DHCP	68	UDP	Out	
DNS	53	TCP/UDP	In/Out	Other cluster nodes and DNS server
Flow Telemetry	5640-5671	UDP	In	In-band of switches Used to receive flow telemetry from fabrics
GRPC (Telemetry)	50051	TCP	In	Information related to multicast flows for IP Fabric for Media deployments as well as PTP for general LAN deployments is streamed out using software telemetry to a persistent IP address associated with a Nexus Dashboard GRPC receiver service pod.
HTTP (PnP)	9666	TCP	In	Cisco Plug and Play (PnP) for Catalyst devices is accomplished using Nexus Dashboard HTTP port 9666 and HTTPS port 9667. HTTP on port 9666 is used to send CA certificate bundle to devices to prime the device for HTTPS mode and actual PnP happens over HTTPS on port 9667 afterwards. PnP service, as with POAP, runs on a persistent IP address that is associated with either the management or data subnet. The persistent IP subnet is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings.

Service	Port	Protocol	Direction (In/Out)	Connection
HTTP (POAP)	80	TCP	In	Only used for device zero-touch provisioning using POAP, where devices can send (limited jailed write-only access to Nexus Dashboard) basic inventory information to Nexus Dashboard to start secure POAP communication. Nexus Dashboard Bootstrap or POAP can be configured for TFTP or HTTP/HTTPS. The SCP-POAP service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings.
HTTPS	443	TCP	Out	In-band of switches and APIC and NX-OS fabrics
HTTPS/HTTP (NX-API)	443/80	TCP	Out	NX-API HTTPS/HTTP client connects to device NX-API server on port 443/80, which is also configurable. NX-API is an optional feature, used by limited set of Nexus Dashboard functions.
HTTPS (PnP)	9667	TCP	In	Cisco Plug and Play (PnP) for Catalyst devices is accomplished using Nexus Dashboard HTTP port 9666 and HTTPS port 9667. HTTP on port 9666 is used to send CA certificate bundle to devices to prime the device for HTTPS mode and actual PnP happens over HTTPS on port 9667 afterwards. PnP service, as with POAP, runs on a persistent IP address that is associated with either the management or data subnet. The persistent IP subnet is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings.

Service	Port	Protocol	Direction (In /Out)	Connection
HTTPS (POAP)	443	TCP	In	Secure POAP is accomplished using the Nexus Dashboard HTTPS Server on port 443. The HTTPS server is bound to the SCP-POAP service and uses the same persistent IP address assigned to that pod. The SCP-POAP service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings.
HTTPS (vCenter, Kubernetes, OpenStack, Discovery)	443	TCP	Out	Nexus Dashboard provides an integrated host and physical network topology view by correlating the information obtained from registered VMM domains, such as VMware vCenter or OpenStack, as well as container orchestrators, such as Kubernetes. This is an optional feature
ICMP	ICMP	ICMP	In/Out	Other cluster nodes, default gateway
Infra-Service	3379 3380 4430 8989 9090 9969 9979 9989 15223 30002-30006 30009-30010 30012 30014-30021 30025 30027 30028	TCP	In/Out	Other cluster nodes

Service	Port	Protocol	Direction (In/Out)	Connection
Infra-Service	30016 30017	TCP/UDP	In/Out	Other cluster nodes
Infra-Service	30019	UDP	In/Out	Other cluster nodes
Infra-Service	30500-30600	TCP/UDP	In/Out	Other cluster nodes
Kafka	30001	TCP	In/Out	In-band IP of switches and APIC/Controller
KMS	9989	TCP	In/Out	Other cluster nodes and ACI fabrics
NFSv3	111	TCP/UDP	In/Out	Remote NFS server
NFSv3	608	UDP	In/Out	Remote NFS server
NFSv3	2049	TCP	In/Out	Remote NFS server
NX-API	8443	TCP	In/Out	Used by Cisco MDS 9000 Series switches with NX-OS release 9.x and later for performance monitoring.
SCP	22	TCP	In/Out	SCP is used by various features to transfer files between devices and Nexus Dashboard, such as for archiving backup files to remote server.. The Nexus Dashboard SCP service serves as the SCP server for both downloads and uploads. SCP is also used by the POAP client on the devices to download POAP-related files. The SCP-POAP service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings.
SCP	22	TCP	Out	Transport tech-support file from persistent IP address of Nexus Dashboard POAP-SCP pod to a separate ND cluster running telemetry. The SCP-POAP service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings

Service	Port	Protocol	Direction (In /Out)	Connection
SMTP	25	TCP	Out	You can configure the SMTP port on the Admin > Server Settings > General page. This is an optional feature.
SNMP	161	TCP/UDP	Out	SNMP traffic from Nexus Dashboard to devices.
SNMP Trap	2162	UDP	In	SNMP traps from devices to Nexus Dashboard are sent out toward the persistent IP address associated with the SNMP-Trap/Syslog service pod. The SNMP-Trap-Syslog service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings,
SSH	22	TCP	Out	UI, in-band of switches, and APIC
SSH	1022	TCP/UDP	In/Out	Other cluster nodes
SW Telemetry	5695 30000 57500 30570	TCP	In/Out	Other cluster nodes Used to collect various telemetry information from fabrics Port 57500 is needed between switches and Nexus Dashboard for telemetry and NX-OS based switches
Tech support	2022	TCP	In	This is used to upload tech support from fabric to cluster.
TFTP (POAP)	69	UDP	In	Only used for device zero-touch provisioning using POAP, where devices can send (limited jailed write-only access to Nexus Dashboard) basic inventory information to Nexus Dashboard to start secure POAP communication. Nexus Dashboard bootstrap or POAP can be configured for TFTP or HTTP/HTTPS. The SCP-POAP service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings.

Service	Port	Protocol	Direction (In/Out)	Connection
VXLAN	4789	UDP	In/Out	Other cluster nodes

Prerequisites for SAN deployments

Network prerequisites for SAN deployments

These network prerequisites apply for SAN deployments:

- In SAN deployments, management and data networks can use the same subnet.
- Persistent IP addresses are supported on the data network only with Layer 2 adjacent and Layer 3 adjacent with BGP configured. However, if you configure the management and data networks to use the same subnet, then you cannot use Layer 3 adjacency because you must configure BGP during the bootstrap process for the data network Layer 3 adjacency, but the management network doesn't support Layer 3 adjacency with BGP configured.

In this situation, you can:

- Use Layer 2 adjacency instead if you want to configure the management and data networks to use the same subnet, or
- Use different subnets for the management and data networks, where you can configure Layer 2 adjacency on the management network and Layer 3 adjacency with BGP configured on the data network
- You must allocate some number of persistent IP addresses depending on your configuration. For more information about persistent IP addresses, see [Nexus Dashboard persistent IP addresses, on page 40](#).

Communication ports for SAN deployments

Nexus Dashboard uses TLS or mTLS with encryption to protect data privacy and integrity while in transit.

This table lists the management network communication ports for SAN deployments.

Table 5: Management network communication ports for SAN deployments

Service	Port	Protocol	Direction In—toward the cluster Out—from the cluster toward the fabric or outside world	Connection
DNS	53	TCP/UDP	Out	DNS server

Service	Port	Protocol	Direction In—toward the cluster Out—from the cluster toward the fabric or outside world	Connection
GRPC (Telemetry)	33000	TCP	In	SAN Telemetry Server which receives SAN data (such as storage, hosts, flows, and so on) over GRPC transport tied to Nexus Dashboard persistent IP address.
HTTP	80	TCP	Out	Internet/proxy
HTTPS	443	TCP	In/Out	UI, other clusters (for multi-cluster connectivity), fabrics, Internet/proxy
HTTPS (vCenter, Kubernetes, OpenStack, Discovery)	443	TCP	Out	Nexus Dashboard provides an integrated host and physical network topology view by correlating the information obtained from registered VMM domains, such as VMware vCenter or OpenStack, as well as container orchestrators, such as Kubernetes. This is an optional feature
ICMP	ICMP	ICMP	In/Out	Other cluster nodes, CIMC, default gateway
Infra-Service	30012 30021 30500-30600	TCP/UDP	In/Out	Other cluster nodes
KMS	9880	TCP	In/Out	Other cluster nodes and ACI fabrics
LDAP	389 636	TCP	Out	LDAP server
NTP	123	UDP	Out	NTP server
NX-API	8443	TCP	In/Out	Used by Cisco MDS 9000 Series switches with NX-OS release 9.x and later for performance monitoring.
Radius	1812	TCP	Out	Radius server

Service	Port	Protocol	Direction In—toward the cluster Out—from the cluster toward the fabric or outside world	Connection
SCP	22	TCP	In/Out	SCP is used by various features to transfer files between devices and Nexus Dashboard, such as for archiving backup files to remote server.. The Nexus Dashboard SCP service serves as the SCP server for both downloads and uploads. SCP is also used by the POAP client on the devices to download POAP-related files. The SCP-POAP service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings.
SCP	22	TCP	Out	Transport tech-support file from persistent IP address of Nexus Dashboard POAP-SCP pod to a separate ND cluster running telemetry. The SCP-POAP service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings
SMTP	25	TCP	Out	You can configure the SMTP port on the Admin > Server Settings > General page. This is an optional feature.
SNMP	161	TCP/UDP	Out	SNMP traffic from Nexus Dashboard to devices.

Service	Port	Protocol	Direction In—toward the cluster Out—from the cluster toward the fabric or outside world	Connection
SNMP Trap	2162	UDP	In	SNMP traps from devices to Nexus Dashboard are sent out toward the persistent IP address associated with the SNMP-Trap/Syslog service pod. The SNMP-Trap-Syslog service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings,
SSH	22	TCP	In/Out	CLI and CIMC of the cluster nodes
Syslog	514	UDP	In	When Nexus Dashboard is configured as a Syslog server, Syslogs from the devices are sent out toward the persistent IP address associated with the SNMP-Trap/Syslog service pod. The SNMP-Trap-Syslog service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings.
TACACS	49	TCP	Out	TACACS server

This table lists the data network communication ports for SAN deployments.

Table 6: Data network communication ports for SAN deployments

Service	Port	Protocol	Direction In—toward the cluster Out—from the cluster toward the fabric or outside world	Connection
DNS	53	TCP/UDP	In/Out	Other cluster nodes and DNS server
GRPC (Telemetry)	33000	TCP	In	SAN Telemetry Server which receives SAN data (such as storage, hosts, flows, and so on) over GRPC transport tied to Nexus Dashboard persistent IP address.
HTTPS	443	TCP	Out	In-band of switches and APIC and NX-OS fabrics
HTTPS (vCenter, Kubernetes, OpenStack, Discovery)	443	TCP	Out	Nexus Dashboard provides an integrated host and physical network topology view by correlating the information obtained from registered VMM domains, such as VMware vCenter or OpenStack, as well as container orchestrators, such as Kubernetes. This is an optional feature
ICMP	ICMP	ICMP	In/Out	Other cluster nodes, default gateway

Service	Port	Protocol	Direction In—toward the cluster Out—from the cluster toward the fabric or outside world	Connection
Infra-Service	3379 3380 8989 9090 9969 9979 9989 15223 30002-30006 30009-30010 30012 30014-30015 30018-30019 30025 30027	TCP	In/Out	Other cluster nodes
Infra-Service	30016 30017	TCP/UDP	In/Out	Other cluster nodes
Infra-Service	30019	UDP	In/Out	Other cluster nodes
Infra-Service	30500-30600	TCP/UDP	In/Out	Other cluster nodes
KMS	9880	TCP	In/Out	Other cluster nodes and ACI fabrics
NFSv3	111	TCP/UDP	In/Out	Remote NFS server
NFSv3	608	UDP	In/Out	Remote NFS server
NFSv3	2049	TCP	In/Out	Remote NFS server
NX-API	8443	TCP	In/Out	Used by Cisco MDS 9000 Series switches with NX-OS release 9.x and later for performance monitoring.

Service	Port	Protocol	Direction In—toward the cluster Out—from the cluster toward the fabric or outside world	Connection
SCP	22	TCP	In/Out	SCP is used by various features to transfer files between devices and Nexus Dashboard, such as for archiving backup files to remote server.. The Nexus Dashboard SCP service serves as the SCP server for both downloads and uploads. SCP is also used by the POAP client on the devices to download POAP-related files. The SCP-POAP service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings.
SCP	22	TCP	Out	Transport tech-support file from persistent IP address of Nexus Dashboard POAP-SCP pod to a separate ND cluster running telemetry. The SCP-POAP service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings
SMTP	25	TCP	Out	You can configure the SMTP port on the Admin > Server Settings > General page. This is an optional feature.
SNMP	161	TCP/UDP	Out	SNMP traffic from Nexus Dashboard to devices.

Service	Port	Protocol	Direction In—toward the cluster Out—from the cluster toward the fabric or outside world	Connection
SNMP Trap	2162	UDP	In	SNMP traps from devices to Nexus Dashboard are sent out toward the persistent IP address associated with the SNMP-Trap/Syslog service pod. The SNMP-Trap-Syslog service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings,
SSH	22	TCP	Out	In-band of switches and APIC
SSH	1022	TCP/UDP	In/Out	Other cluster nodes
Syslog	514	UDP	In	When Nexus Dashboard is configured as a Syslog server, Syslogs from the devices are sent out toward the persistent IP address associated with the SNMP-Trap/Syslog service pod. The SNMP-Trap-Syslog service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings.
VXLAN	4789	UDP	In/Out	Other cluster nodes

This table lists the ports that the Nexus Dashboard SAN deployments on single-node clusters require.

Table 7: Nexus Dashboard ports for SAN deployments on single-node clusters

Service	Port	Protocol	Direction In—toward the cluster Out—from the cluster toward the fabric or outside world	Connection (Applies to both LAN and SAN deployments, unless stated otherwise)
GRPC (Telemetry)	33000	TCP	In	SAN Telemetry Server which receives SAN data (such as storage, hosts, flows, and so on) over GRPC transport tied to Nexus Dashboard persistent IP address.
HTTPS (vCenter, Kubernetes, OpenStack, Discovery)	443	TCP	Out	Nexus Dashboard provides an integrated host and physical network topology view by correlating the information obtained from registered VMM domains, such as VMware vCenter or OpenStack, as well as container orchestrators, such as Kubernetes. This is an optional feature.
SCP	22	TCP	In/Out	SCP is used by various features to transfer files between devices and Nexus Dashboard, such as for archiving backup files to remote server.. The Nexus Dashboard SCP service serves as the SCP server for both downloads and uploads. SCP is also used by the POAP client on the devices to download POAP-related files. The SCP-POAP service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings.
SCP	22	TCP	Out	Transport tech-support file from persistent IP address of Nexus Dashboard POAP-SCP pod to a separate ND cluster running telemetry. The SCP-POAP service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings

Service	Port	Protocol	Direction In—toward the cluster Out—from the cluster toward the fabric or outside world	Connection (Applies to both LAN and SAN deployments, unless stated otherwise)
SMTP	25	TCP	Out	You can configure the SMTP port on the Admin > Server Settings > General page. This is an optional feature.
SNMP	161	TCP/UDP	Out	SNMP traffic from Nexus Dashboard to devices.
SNMP Trap	2162	UDP	In	SNMP traps from devices to Nexus Dashboard are sent out toward the persistent IP address associated with the SNMP-Trap/Syslog service pod. The SNMP-Trap-Syslog service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings,
SSH	22	TCP	Out	SSH is a basic mechanism for accessing devices.
Syslog	514	UDP	In	When Nexus Dashboard is configured as a Syslog server, Syslogs from the devices are sent out toward the persistent IP address associated with the SNMP-Trap/Syslog service pod. The SNMP-Trap-Syslog service in Nexus Dashboard has a persistent IP address that is associated with either the management or data subnet. This is controlled by the LAN Device Management Connectivity setting in the Nexus Dashboard server settings.

Nexus Dashboard persistent IP addresses

Persistent IP addresses, also known as external service IP addresses, are IP addresses that are used for various controller and telemetry functions within the Nexus Dashboard cluster. The word “persistent” is used because while the service may move between different Nexus Dashboard nodes in the event of a node or pod failure, the IP address for the service being referred by the switches in the fabric is preserved. This ensures that no

configuration updates to the switches are needed in case of a Nexus Dashboard-related failure event. Persistent IP addresses can be programmed in both the management and data subnets depending on the features deployed.

You can view the configured persistent IP addresses on your Nexus Dashboard by navigating to:

Admin > System Settings > General

Locate the **External Pools** tile and click **View all** at the bottom left area of the **External Pools** tile to view the configured persistent IP addresses on your Nexus Dashboard.

Persistent IP address updates in release 4.x

This section provides information on the changes in Nexus Dashboard release 4.x for the persistent IP addresses. It also explains how to make certain updates in the number of persistent IP addresses before proceeding to upgrade to Nexus Dashboard release 4.x.

- Reduction of number of IP addresses needed

Starting with release 4.1.1, some services that needed exclusive IP addresses in previous Nexus Dashboard releases were merged with others. For example, on a per Nexus Dashboard node basis, software telemetry, flow telemetry and IPFM telemetry collectors on the data network have been merged in such a way that one IP address per collector service is sufficient to serve all three functions.

- LAN Device Connectivity

You can set the type of LAN device connectivity (Data or Management) under **Admin > System Settings > Fabric management > Advanced settings > Admin > LAN Device Management Connectivity**.

Prior to release 4.1.1, the default setting for LAN Device connectivity was Management. Beginning with release 4.1.1, this default has been changed to Data. However, when upgrading from Nexus Dashboard release 3.2.x to 4.x, any user-configured setting for LAN device connectivity is preserved.

- Layer-3 Persistent IP Support for Telemetry

Starting with Nexus Dashboard release 4.1.1, telemetry-collector-X persistent IPs are supported in layer-3 adjacent Nexus Dashboard clusters. See below for Layer 3 BGP deployment details.

The number of persistent IP addresses and how they are mapped to services has changed in Nexus Dashboard release 4.x. The following services consume persistent IP addresses on Nexus Dashboard:

LAN deployments:

- Telemetry collector-x: There is a requirement for 1 persistent IP address for a 1-node cluster, or 3 persistent IP addresses for a 3-node or larger cluster, on the data network.
- SNMP trap and syslog receiver: 1 persistent IP address on the data network if the LAN device connectivity type is set to Data or 1 persistent IP address on the management network if the LAN device connectivity type is set to Management.
- Switch bootstrap service (POAP/PnP): 1 persistent IP address on the data network if the LAN device connectivity type is set to Data or 1 persistent IP address on the management network if the LAN device connectivity type is set to Management.
- (Optional) Endpoint Locator (EPL): 1 persistent IP address on the data network for each fabric where EPL is enabled. The EPL feature can be enabled for up to 4 fabrics in a given Nexus Dashboard cluster.
- (Optional) IPFM (IP Fabric for Media) telemetry collector-x: If LAN device connectivity is set to Data, then no additional persistent IPs are required. However if LAN device connectivity type is set to

Management, there is a requirement for 1 persistent IP address for a 1-node cluster, or 3 persistent IP addresses for a 3-node or larger cluster, on the management network.

SAN deployments:

- SNMP trap and syslog receiver: 1 persistent IP address on the data network.
- Switch bootstrap service: 1 persistent IP address on the data network.
- (Optional) SAN Insights receiver-x: There is a requirement for 1 persistent IP address for a 1-node cluster, or 3 persistent IP addresses for a 3-node or larger cluster, on the data network.

For an IPv4-only Nexus Dashboard cluster deployment, each service listed above will consume 1 persistent IPv4 address. For an IPv6-only Nexus Dashboard cluster deployment, each service will consume 1 persistent IPv6 address. For a dual-stack Nexus Dashboard cluster deployment, each service needing a persistent IP will consume an IPv4 address and an IPv6 address.

Fresh installation or upgrade

The total number of persistent IP addresses that you will need won't necessarily change based on whether you are performing a fresh installation or an upgrade. In addition, for a fresh installation of Nexus Dashboard (greenfield deployment), you must configure persistent IP addresses only on the data network. You can change the LAN device connectivity type from Data to Management after the cluster installation is complete.

As mentioned earlier, there is a change in the default setting in Nexus Dashboard 4.x compared to previous Nexus Dashboard releases. In Nexus Dashboard 4.x, the default LAN Device Management Connectivity is set to Data, while in earlier releases it was set to Management. As part of the effort towards delivering the unified Nexus Dashboard, the goal is to make the recommended best practice deployment as easy as possible. The reachability from Nexus Dashboard to and from switches is recommended to be over the Nexus Dashboard data interface. The Nexus Dashboard management interface should primarily be used for UI/API access and reachability for AAA, DNS, Proxy, NTP, Intersight, and so on. Finally, note that the connectivity setting set by the user is preserved when doing an inline upgrade from Nexus Dashboard release 3.2.x to Nexus Dashboard release 4.x.

Additional considerations to keep in mind

Along with the factors listed above, there are several additional considerations that you should keep in mind regarding persistent IP addresses:

- Nexus Dashboard deployment mode:
 - Layer 2: Here the Nexus Dashboard nodes within the cluster are layer-2 adjacent. This means that all Nexus Dashboard nodes share the same management and data subnet respectively. Persistent IP addresses need to be on the same network as the data network or management network.
 - Layer 3 BGP: In this mode, the Nexus Dashboard nodes within the cluster are layer-3 adjacent. In other words, unique management and data subnets are associated with each Nexus Dashboard node in the cluster. There needs to be IP reachability between the nodes to form the cluster. Persistent IP addresses cannot be from a subnet that belongs to any of the Nexus Dashboard nodes' Data or Management interface subnets. In this case, LAN Device Management Connectivity must be set to Data and cannot be changed.

Updates to mapping

The mapping for persistent IP addresses has been updated to show correct service names.

Older	Newer
cisco-nir-collectorpersistent1-service	Telemetry collector-1
cisco-nir-collectorpersistent2-service	Telemetry collector-2
cisco-nir-collectorpersistent3-service	Telemetry collector-3
cisco-ndfc-dcnm-poap-data-http-ssh	Switch Bootstrap server
cisco-ndfc-dcnm-poap-mgmt-http-ssh	Switch Bootstrap server
cisco-ndfc-dcnm-syslog-trap-data	SNMP trap and syslog receiver
cisco-ndfc-dcnm-syslog-trap-mgmt	SNMP trap and syslog receiver
cisco-ndfc-pmn-telemetry-mgmt-worker-0	IPFM telemetry collector-1
cisco-ndfc-pmn-telemetry-mgmt-worker-1	IPFM telemetry collector-2
cisco-ndfc-pmn-telemetry-mgmt-worker-2	IPFM telemetry collector-3
cisco-ndfc-dcnm-san-insight-receiver-1	SAN Insights receiver-1
cisco-ndfc-dcnm-san-insight-receiver-2	SAN Insights receiver-2
cisco-ndfc-dcnm-san-insight-receiver-3	SAN Insights receiver-3

Determining the total number of persistent IP addresses that you will need

All the factors listed above come into play when you are trying to determine the total number of persistent IP addresses that you will need and which network they come from. Make sure you review your final Nexus Dashboard deployment configuration to verify that you have enough persistent IP addresses in the proper subnet range for your deployment, and that you have additional persistent IP addresses if necessary, depending on the type of LAN device connectivity that you set and for services that you might enable, such as Endpoint Locator (EPL).

Following is an example scenario that demonstrates how persistent IP addresses are used:

Fresh installation:

First, at cluster bringup, you will need a certain number of persistent IP addresses on the data network, based on the cluster size, as mentioned earlier:

- **1-node cluster, with physical or virtual nodes:** Minimum of 3 persistent IP addresses needed on the data network
- **3-node or larger cluster, with physical or virtual nodes:** Minimum of 5 persistent IP addresses needed on the data network



Note

These values are valid for either IPv4 or IPv6, but double the number for dual-stack IPv4 and IPv6. For example, for a 3-node or larger cluster, a minimum of 10 persistent IP addresses would be needed on the data network for dual-stack IPv4 and IPv6 (5 IPv4 and 5 IPv6 persistent IP addresses).

After bootstrapping, you may need to add additional persistent IP addresses as needed, depending on these scenarios:

- If you leave the LAN device connectivity type set to **Data**, you won't need any additional persistent IP addresses unless you enable the Endpoint Locator (EPL) feature, which requires 1 additional persistent IP address on the data network per fabric where EPL is enabled.
- If you change the LAN device connectivity type from **Data** to **Management**:
 - You will need 2 additional persistent IP addresses on the management network for Syslog/SNMP trap and switch bootstrap functionality.
 - (Optional) If you want to enable Endpoint Locator (EPL), you will need 1 persistent IP address on the data network per fabric where EPL is enabled.
 - (Optional) If you want IP Fabric for Media (IPFM) fabrics, you will need 1 persistent IP address for a 1-node cluster, or 3 persistent IP addresses for a 3-node or larger cluster, on the management network.

Table 8: Persistent IP requirements: Fresh installation of 4.x

Number of ND nodes	LAN Device Connectivity	Mandatory persistent IP addresses	Optional persistent IP addresses	Other common persistent IP addresses
1	Data ¹	3 in data network	N/A	1 in data network per fabric where EPL is enabled
	Management	2 in management network 1 in data network	1 in management network for IPFM fabrics	
3 or more	Data ¹	5 in data network	N/A	
	Management	2 in management network 3 in data network	3 in management network for IPFM fabrics	

¹ Indicates default option set during ND bootstrap process

Upgrade:

Now assume that you want to upgrade from Nexus Dashboard 3.2.x to 4.x. The number of persistent IP addresses that you will need in Nexus Dashboard 4.x will vary, depending on the services that you were running and how they were configured on Nexus Dashboard 3.2.x, and the size of your cluster in Nexus Dashboard 4.x. Note that the LAN Device Management Connectivity that you set in the Nexus Dashboard 3.2.x release is preserved as-is when performing an inline upgrade to the Nexus Dashboard 4.x release.

- If you had only **NDFC** running in your Nexus Dashboard 3.2.x system, and
 - If you had **Data** set as the type of LAN device connectivity in Nexus Dashboard 3.2.x, and
 - You have a 1-node cluster that you want to upgrade to Nexus Dashboard 4.x, then you'll need 3 persistent IP addresses on the data network.
 - You have a 3-node or larger cluster that you want to upgrade to Nexus Dashboard 4.x, then you'll need 5 persistent IP addresses on the data network.

- If you had **Management** set as the type of LAN device connectivity in Nexus Dashboard 3.2.x, and
 - You have a 1-node cluster that you want to upgrade to in Nexus Dashboard 4.x, then you would already have either 2 or 3 persistent IP addresses (additional IP is required if IPFM/PTP feature was enabled) in the management network. In addition, you will need 1 persistent IP address in the data network, otherwise the upgrade to 4.x will fail during the pre-upgrade validation step.
 - You have a 3-node or larger cluster that you want to upgrade to in Nexus Dashboard 4.x, then you would already have either 2 or 5 persistent IP addresses (3 additional IPs are required if IPFM/PTP feature was enabled) in the management network. You will need to configure 3 persistent IP addresses in the data network; only then will the upgrade to 4.x proceed.
- If you had only **NDI** running in your Nexus Dashboard 3.2.x system, and
 - You have a 1-node cluster that you want to upgrade to Nexus Dashboard 4.x, then you would already have configured 4 persistent IP addresses in the data network. After the upgrade to 4.x, only 3 persistent IP addresses will be in use. The remaining can be reclaimed.
 - You have a 3-node or larger cluster that you want to upgrade to Nexus Dashboard 4.x, then you would already have configured 8 persistent IP addresses in the data network and 2 additional data IPs for support of standalone NX-OS deployments. After the upgrade to 4.x, only 5 of these data IP addresses will be in use. The remaining can be reclaimed.
- If you had only **NDO** running in your Nexus Dashboard 3.2.x system, then you did not have any persistent IP addresses in your Nexus Dashboard 3.2.x system. When you upgrade to Nexus Dashboard 4.x, if you have a 3-node cluster that you want to upgrade to Nexus Dashboard 4.x, then you'll need 5 persistent IP addresses on the data network before the upgrade can proceed.
- If you had an **NDO** and **NDI** deployment mode in your Nexus Dashboard 3.2.x system, and you have a 3-node or larger cluster that you want to upgrade to Nexus Dashboard 4.x, then you would already have configured 8 persistent IP addresses in the data network. After the upgrade to 4.x, only 5 of these data persistent IP addresses will be in use. The remaining persistent IP addresses can be reclaimed.
- If you had only **NDFC** and **NDI** deployment mode in your Nexus Dashboard 3.2.x system, and you have a 3-node or larger physical ND cluster that you want to upgrade to Nexus Dashboard 4.x, then there were two options based on the LAN Device Management Connectivity setting:
 - If you had **Management** set as the type of LAN device connectivity in Nexus Dashboard 3.2.x, you would already have configured 8 persistent IP addresses in the data network for NDI and 2 persistent IP addresses in the management network for NDFC. After the upgrade to 4.x, the 2 persistent IP addresses in the management subnet will be used along with only 3 of the data persistent IP addresses. The remaining persistent IP addresses can be reclaimed.
 - If you had **Data** set as the type of LAN device connectivity in Nexus Dashboard 3.2.x, you would already have configured 8 persistent IP addresses in the data network for NDI and 2 additional data persistent IP addresses for NDFC. After the upgrade to 4.x, only 5 of these data persistent IP addresses will be in use. The remaining persistent IP addresses can be reclaimed.

Table 9: Persistent IP requirements: Upgrade from 3.2.x to 4.x

ND 3.2.x deployment mode	Number of ND nodes	LAN Device Connectivity	ND 3.2.x persistent IP address requirement	ND 4.x persistent IP address requirement
NDFC	1	Data	2 in data network, plus 1 in data network if IPFM/PTP is enabled	3 in data network
		Management	2 in management network, plus 1 in management network if IPFM/PTP is enabled	2 in management network, plus 1 in management network for IPFM fabrics 1 in data network
NDFC	3 or more	Data	2 in data network, plus 3 in data network if IPFM/PTP is enabled	5 in data network
		Management	2 in management network, plus 3 in management network if IPFM/PTP is enabled	2 in management network, plus 3 in management network for IPFM fabrics 3 in data network
NDFC + NDI	3 physical	Data	10 in data network	5 in data network
		Management	2 in management network 8 in data network	2 in management network, plus 3 in management network for IPFM fabrics 3 in data network
NDI	1	N/A	3 in data network	3 in data network
NDI	3 or more	N/A	10 in data network	5 in data network
NDO	3	N/A	None	5 in data network
NDO + NDI	3 or more	N/A	8 in data network	5 in data network



Note EPL persistent IP address requirements remain the same in release 4.x as they were in release 3.2.x.

BGP configuration and persistent IP addresses

Some prior releases of Nexus Dashboard allowed you to configure one or more persistent IP addresses that require retaining the same IP addresses even in case they are relocated to a different Nexus Dashboard node. However, in those releases, the persistent IP addresses had to be part of the management and data subnets and the feature could be enabled only if all nodes in the cluster were part of the same Layer 3 network. Here, the services used Layer 2 mechanisms such as gratuitous ARP or neighbor discovery to advertise the persistent IP addresses within its Layer 3 network.

While that is still supported, this release also allows you to configure the persistent IP addresses feature even if you deploy the cluster nodes in different Layer 3 networks. In this case, the persistent IP addresses are advertised out of each node's data links using BGP, which we refer to as "Layer 3 mode". The IP addresses must also be part of a subnet that is not overlapping with any of the nodes' management or data subnets. If the persistent IP addresses are outside the data and management networks, this feature will operate in Layer 3 mode by default; if the IP addresses are part of those networks, the feature will operate in Layer 2 mode. BGP can be enabled during cluster deployment or from the Nexus Dashboard GUI after the cluster is up and running.

If you plan to enable BGP and use the persistent IP address functionality, you must:

- Ensure that the peer routers exchange the advertised persistent IP addresses between the nodes' Layer 3 networks.
- For data network Layer 3 adjacency, you must configure BGP during the bootstrap process. BGP does not support management network Layer 3 adjacency.
- Ensure that the persistent IP addresses you allocate do not overlap with any of the nodes' management or data subnets.

Round trip time requirements

Connectivity between the nodes is required on both networks with additional round trip time (RTT) requirements, as listed in this table.

Table 10: Cluster round trip time requirements

Connectivity	Maximum RTT
Between nodes within the same Nexus Dashboard cluster (intra-cluster connectivity)	50 ms
Between nodes in one cluster and nodes in a different cluster if the clusters are connected using multi-cluster connectivity (inter-cluster connectivity) For more information about multi-cluster connectivity, see <i>Connecting Clusters</i> .	500 ms
Between external DNS servers and the Nexus Dashboard cluster	5 seconds
Between cluster nodes and switches	150 ms

Prerequisites for Nexus Dashboard storage (virtual form factors)

This section of the document provides guidance for designing storage infrastructure for Cisco Nexus Dashboard deployments (VMware, KVM, Nutanix, and so on).

The Nexus Dashboard offers features that process high-frequency telemetry and flow data from network devices. As a result, storage plays a critical role in ensuring:

- Reliable data ingestion
- Accurate analytics and insights
- Consistent system performance at scale

This section outlines the recommended storage characteristics, performance expectations, and architectural considerations for production environments.

Key workload characteristics

These are the key workload characteristics:

- Write-intensive ingestion with sustained throughput requirements
- Mixed I/O patterns:
 - Sequential writes (data ingestion)
 - Random reads/writes (queries and analytics)
- Latency-sensitive operations, particularly during ingestion

Baseline storage requirements

These are the baseline storage requirements:

Capacity

- Storage capacity: Two storage disks are required:
 - One disk of size 50G used as boot disk
 - Another disk to store data
- The size of data depends on platform flavor. For example: Virtual-data requires 3 TB of disk.

Performance

Designing storage for Nexus Dashboard requires careful attention to performance, not just capacity. You should ensure:

- Sync latency is less than 10ms.
- The actual throughput and IOPS requirement depends on scale and can vary from:
 - IOPS: 1.5K io/s to 80K io/s sustained
 - Throughput: 10 MB/s to 1.5 GB/s sustained

You can also verify disk I/O latency using the `acs diag diskio /dev/sdb` command from Nexus Dashboard.

Recommendations for designing storage

To meet these requirements, these are the recommendations for designing storage for Nexus Dashboard.

Storage media selection and placement

For production deployments, we strongly recommend using NVMe or SSDs. From a placement point of view, direct-attached storage (DAS) typically provides the most predictable results. If you are using RAID controllers, configure in JBOD mode to minimize overhead.

Using remote or shared storage

You can use remote storage solutions (such as SAN, NAS, or software-defined storage platforms) if they consistently meet performance requirements. Note that storage must deliver required IOPS and throughput continuously, not just during short bursts. In many cases, locally attached storage provides the most consistent and predictable performance for Nexus Dashboard workloads.

- In case of shared storage, performance variability may occur due to other workloads and resource contention can introduce latency spikes.
- Therefore, we recommend that you have dedicated disks and not shared along with other nodes within the Nexus Dashboard cluster or with other workloads.

Caching

- From the Nexus Dashboard point of view, every write is flushed via hypervisor; therefore, caching plays a critical role to achieve performance.
- Without a cache, every single 4KB write has to be written all the way to the physical medium. You should always enable cache policy to write back instead of Write Through.
- Note that hardware vendors typically will set the default to Write Through to guarantee data integrity. For example, if the power cord is pulled or the building loses power at that exact second, the data is already on the disk. It is completely safe.
- Modern systems use battery (BBU) and they use cache only if the battery or SuperCap is healthy. Make sure that your disk controller is provisioned to fallback automatically to Write Through if the battery fails or is charging.

Fabric connectivity

These sections describe how to connect your Nexus Dashboard cluster nodes to the management and data networks and how to connect the cluster to your fabrics. For more information on configuring the fabric to enable the in-band telemetry functions, see the following documents:

- [Getting Your Cisco ACI Fabrics Ready for Cisco Nexus Dashboard Insights](#)
- [Deploying Nexus Fabrics with Telemetry on Cisco Nexus Dashboard](#)

For on-premises APIC or NX-OS fabrics, you can connect the Nexus Dashboard cluster in one of these ways:

- The Nexus Dashboard cluster connected to the fabric using a Layer 3 network.
- The Nexus Dashboard nodes connected to the leaf switches as typical hosts.

Connecting using an external Layer 3 network

We recommend connecting the Nexus Dashboard cluster to the fabrics using an external Layer 3 network as it does not tie the cluster to any one fabric and the same communication paths can be established to all fabrics. Specific connectivity depends on the type of applications deployed in the Nexus Dashboard:

- If you are using orchestration to manage Cisco ACI fabrics, you can establish connectivity from either the data interface or the management interface to either the in-band or out-of-band (OOB) interface of each fabric's APIC or both.

If the fabric connectivity is from the Nexus Dashboard's management interface, you must configure specific static routes or ensure that the management interface is part of the same IP subnet of the APIC interfaces.

- If you are using telemetry, you must establish connectivity from the data interface to the in-band network of each fabric and of the APIC.

If you plan to connect the cluster across a Layer 3 network, keep the following in mind:

- For ACI fabrics, you must configure an L3Out and the external EPG for Cisco Nexus Dashboard data network connectivity in the management tenant.

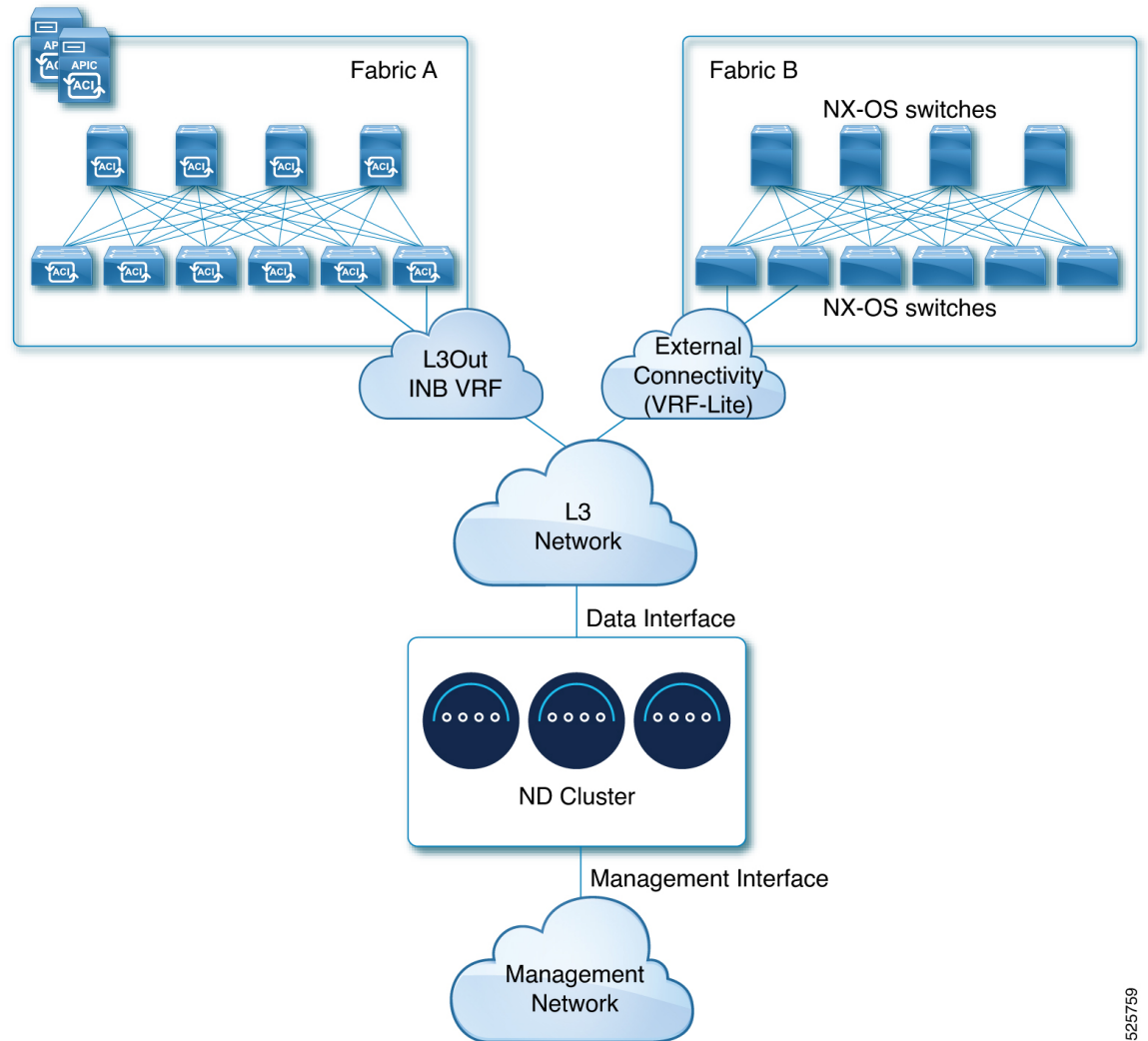
Configuring external connectivity in an ACI fabric is described in the [Cisco APIC Layer 3 Networking Configuration Guide](#).

- If you specify a VLAN ID for your data interface during setup of the cluster, the host port must be configured as `trunk` allowing that VLAN.

However, in most common deployments, you can leave the VLAN ID empty and configure the host port in `access` mode.

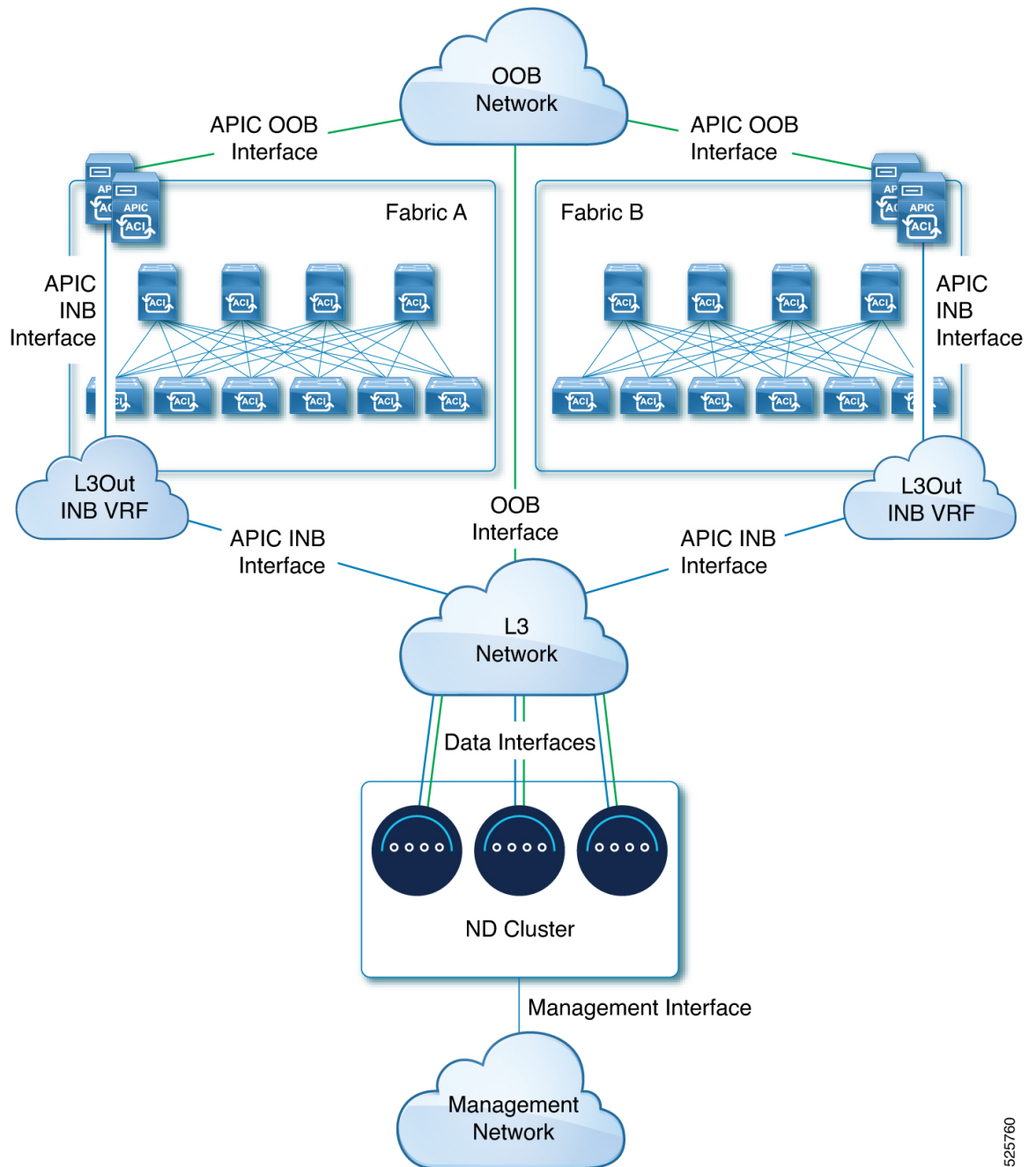
The following two figures show two distinct network connectivity scenarios when connecting the Nexus Dashboard cluster to the fabrics using a Layer 3 network, where the first figure shows a mix of ACI and NX-OS fabrics, and the second figure shows only ACI fabrics.

Figure 1: Connecting using Layer 3 network, With a mix of ACI and NX-OS fabrics



525759

Figure 2: Connecting using Layer 3 network, with ACI fabrics only



Connecting nodes directly to leaf switches

You can also connect the Nexus Dashboard cluster directly to one of the fabrics. This provides easy connectivity between the cluster and in-band management of the fabric, but ties the cluster to the specific fabric and requires reachability to other fabrics to be established through external connectivity. This also makes the cluster dependent on the specific fabric so issues within the fabric may impact Nexus Dashboard connectivity. Like in the previous example, connectivity depends on the type of applications deployed in the Nexus Dashboard:

- If you are using orchestration to manage Cisco ACI fabrics, you can establish connectivity from either the data interface or the management interface to either the in-band or out-of-band (OOB) interface of each fabric's APIC or both.

If the fabric connectivity is from the Nexus Dashboard's management interface, you must configure specific static routes or ensure that the management interface is part of the same IP subnet of the APIC interfaces.

- If you are using telemetry, you can establish connectivity from the data interface to the in-band or out-of-band (OOB) interface of each fabric. However, you must add the route if you establish connectivity from the data interface to the out-of-band interface.

For ACI fabrics, the data interface IP subnet connects to an EPG/ or bridge domain in the fabric and must have a contract established to the local in-band EPG in the management tenant. We recommend deploying the Nexus Dashboard in the management tenant and in-band VRF. Connectivity to other fabrics is established using an L3Out.

If you plan to connect the cluster directly to the leaf switches, keep the following in mind:

- If deploying in VMware ESX or Linux KVM, the host must be connected to the fabric using trunk port.
- If you are connecting Nexus Dashboard directly to NX-OS leaf switches, and those switches are also going to be managed by Nexus Dashboard, the configuration for the data network must be provisioned manually using the CLI so that the cluster can be bootstrapped before adding the switches to a new fabric. Because of this scenario, if you plan to have the fabric to which Nexus Dashboard is connected managed by the same Nexus Dashboard, we do not recommend that you deploy new fabrics in this way.
- If you specify a VLAN ID for your data network during setup of the cluster, the Nexus Dashboard interface and the port on the connected network device must be configured as `trunk`.

However, in most cases we recommend not assigning a VLAN to the data network, in which case you must configure the ports in `access` mode.

- For configurations on the APIC side, following are required configurations:
 - We recommend configuring the bridge domain, subnet, and endpoint group (EPG) for Cisco Nexus Dashboard connectivity in management tenant.
Because the Nexus Dashboard requires connectivity to the in-band EPG in the default in-band VRF in the mgmt tenant, creating the EPG in the management tenant means no route leaking is required.
 - You must create a contract between the fabric's in-band management EPG and Cisco Nexus Dashboard EPG.
 - If several fabrics are monitored on the Nexus Dashboard cluster, L3Out with default route or specific route to other ACI fabric in-band EPG must be provisioned and a contract must be established between the cluster EPG and the L3Out's external EPG.

The following figures show two distinct network connectivity scenarios when connecting the Nexus Dashboard cluster directly to the fabrics' leaf switches. The primary purpose of each depends on the type of application you may be running in your Nexus Dashboard.

The following graphics show these types of connections:

- Connecting directly to ACI fabric
- Connecting directly to NX-OS fabric

- Connecting directly to ACI and NX-OS fabrics

Figure 3: Connecting Directly to ACI Fabric

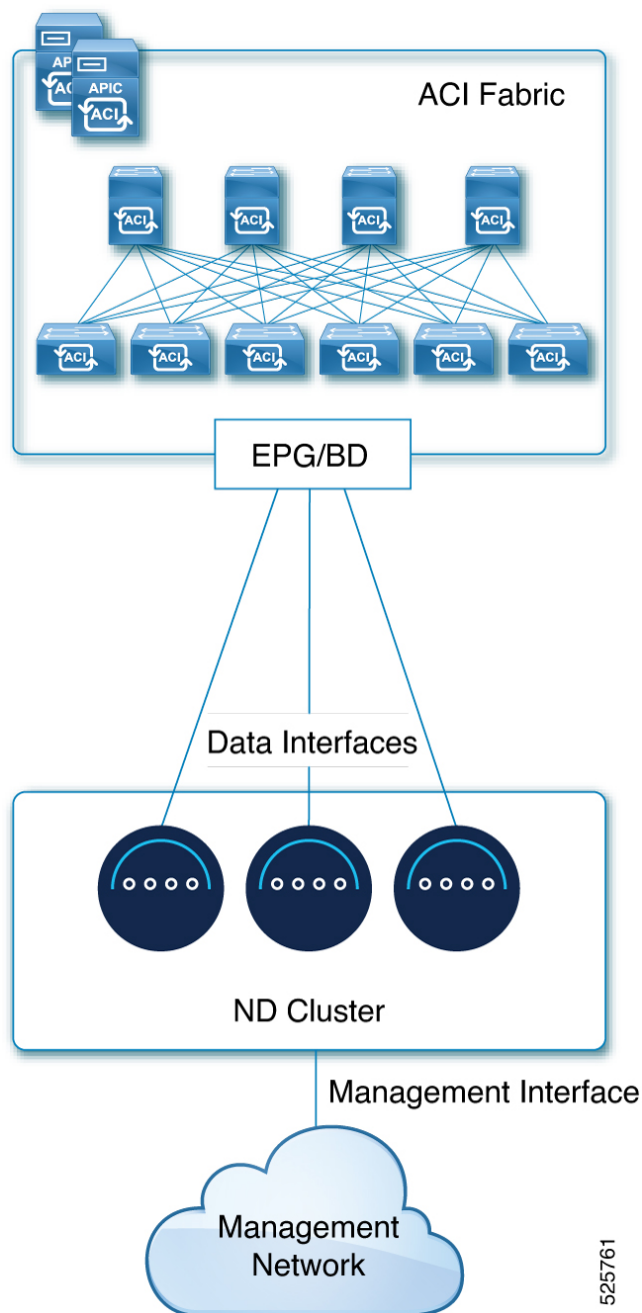


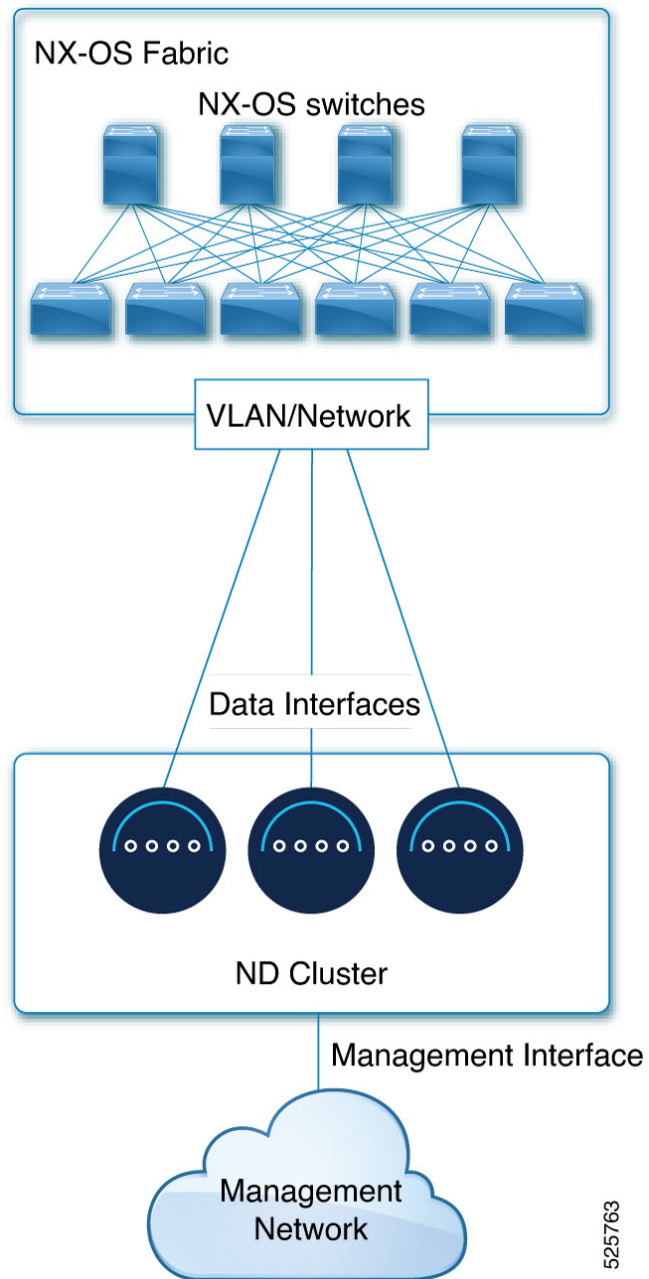
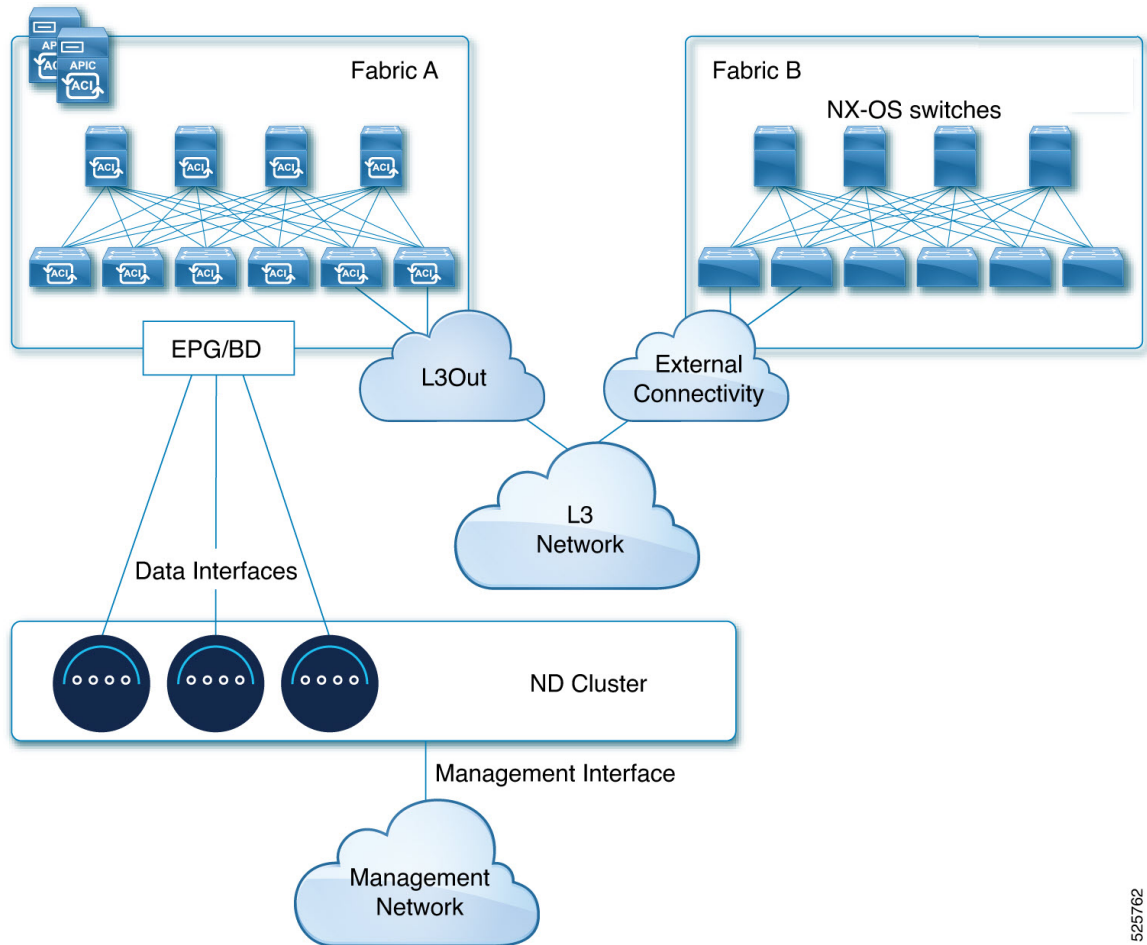
Figure 4: Connecting Directly to NX-OS Fabric

Figure 5: Connecting Directly to ACI and NX-OS Fabrics



525762



PART II

Deploying the Cluster

- [Pre-Installation Checklist, on page 59](#)
- [Deploying as a Physical Appliance, on page 63](#)
- [Deploying in VMware ESX, on page 119](#)
- [Deploying in Linux KVM, on page 147](#)
- [Deploying a Virtual Nexus Dashboard \(vND\) in Nutanix, on page 169](#)
- [Deploying a Virtual Nexus Dashboard \(vND\) in Amazon Web Services \(AWS\), on page 203](#)



CHAPTER 4

Pre-Installation Checklist

- [Pre-installation checklist, on page 59](#)

Pre-installation checklist

Before you proceed with deploying your Nexus Dashboard cluster, prepare this information for easy reference during the process.

Table 11: Cluster details

Parameters	Example	Your entry
Cluster Name	ND-Prod-CL01	
Nexus Dashboard Implementation Type (LAN/SAN)	LAN	
DNS Provider(s)	8.8.8.8	
DNS Search Domain	cisco.com	
NTP Provider(s)	1.1.1.1	
(Optional) NTP Authentication Key	123456789	
(Optional) NTP Authentication ID	100	
(Optional) NTP Authentication Type	MD5	
Proxy Server	192.168.50.1	
(Optional) Proxy Server Ignore Hosts	10.0.0.1	
(Optional) Proxy Username	Proxy-user	

Parameters	Example	Your entry
(Optional) Proxy Password	P@ssword!123	
App Network	172.17.0.1/16 (default) ¹	
Services Network	100.80.0.0/16 (default) ²	
(Optional) App Network IPv6	2001:db8:abcd:0012::0/64	
(Optional) Services Network IPv6	2001:db8:efgh:0012::0/64	

² These are the default values and we do not recommend that you change them. If you want to change the values, see the appropriate chapter in the "Deploying the Cluster" part.



Note You can define all nodes during the initial cluster deployment, including the `secondary` and `standby` nodes. For simplicity, the following tables assume a 3-node base cluster, but if you are deploying a larger cluster, you must also have the details for all additional nodes.

Table 12: Node details

Parameters	Example	Your entry
For physical nodes, CIMC address and login information of the first node	10.196.220.84/24 Username: admin Password: Cisco1234	
For physical nodes, CIMC address and login information of the second node	10.196.220.85/24 Username: admin Password: Cisco1234!	
For physical nodes, CIMC address and login information of the third node	10.196.220.86/24 Username: admin Password: Cisco1234!	
Password used for each node's <code>rescue-user</code> and the initial GUI password. We recommend configuring the same password for all nodes in the cluster.	Welcome2Cisco!	
Management IP of the first node	192.168.11.172/24	
Management Gateway of the first node.	192.168.11.1	
Data Network IP of the first node	192.168.8.172/24	

Parameters	Example	Your entry
Data Network Gateway of the first node	192.168.8.1	
(Optional) Data Network VLAN of the first node (only enter a VLAN if the upstream switchport configuration is “trunk” and the VLAN is added to the trunk allowed list)	101	
(Optional) If you enable BGP, ASN of the first node	63331	
(Optional) If you enable BGP and use IPv6-only deployment, Router ID for the first node in the form of an IPv4 address	1.1.1.1	
(Optional) If you enable BGP, IP addresses of the first node's BGP Peers	200.11.11.2	
(Optional) If you enable BGP, ASNs of the first node's BGP Peers	55555	
Management IP of the second node	192.168.9.173/24	
Management Gateway of the second node.	192.168.9.1	
Data Network IP of the second node	192.168.6.173/24	
Data Network Gateway of the second node	192.168.6.1	
(Optional) Data Network VLAN of the second node	101	
(Optional) If you enable BGP, ASN of the second node	63331	
(Optional) If you enable BGP and use IPv6-only deployment, Router ID for the second node in the form of an IPv4 address	2.2.2.2	
(Optional) If you enable BGP, IP addresses of the second node's BGP Peers	200.12.12.2	

Parameters	Example	Your entry
Cluster Connectivity (L2/BGP)s	L2	
(Optional) If you enable BGP, ASNs of the second node's BGP Peers	55555	
Management IP of the third node	192.168.9.174/24	
Management Gateway of the third node.	192.168.9.1	
Data Network IP of the third node	192.168.6.174/24	
Data Network Gateway of the third node	192.168.6.1	
(Optional) Data Network VLAN of the third node	101	
(Optional) If you enable BGP, ASN of the third node	63331	
(Optional) If you enable BGP and use IPv6-only deployment, Router ID in the form of an IPv4 address	3.3.3.3	
(Optional) If you enable BGP, IP addresses of the third node's BGP Peers	200.13.13.2	
(Optional) If you enable BGP, ASNs of the third node's BGP Peers	55555	

You will also need to program persistent IP addresses during the cluster bringup. For more information, see [Nexus Dashboard persistent IP addresses, on page 40](#).



CHAPTER 5

Deploying as a Physical Appliance

- [Prerequisites and guidelines for deploying Nexus Dashboard as a physical appliance, on page 63](#)
- [Understanding and cabling physical appliances, on page 65](#)
- [Deploy Nexus Dashboard as a physical appliance, on page 110](#)

Prerequisites and guidelines for deploying Nexus Dashboard as a physical appliance

Before you proceed with deploying the Nexus Dashboard cluster, you must:

- Review and complete the prerequisites described in [Prerequisites and Guidelines, on page 11](#).
- Review the *Cisco Nexus Dashboard Release Notes* for any information that can affect your deployment. See the [Cisco Nexus Dashboard documentation landing page](#).
- Ensure you are using the following hardware and the servers are racked and connected as described in [Cisco Nexus Dashboard Hardware Setup Guide](#) specific to the model of server you have.

The physical appliance form factor is supported only on these versions of the original Cisco Nexus Dashboard platform hardware:

- SE-NODE-G2 (UCS-C220-M5). The product ID of the 3-node cluster chassis is SE-CL-L3.
- ND-NODE-L4 (UCS-C225-M6). The product ID of the 3-node cluster chassis is ND-CLUSTER-L4.
- ND-NODE-G5S (UCS-C225-M8). The product ID of the 3-node cluster chassis is ND-CLUSTERG5S.
- ND-NODE-G5L (UCS-C225-M8). The product ID of the 3-node cluster chassis is ND-CLUSTERG5L.



Note This hardware only supports Cisco Nexus Dashboard software. If any other operating system is installed, the node can no longer be used as a Cisco Nexus Dashboard node.

- Ensure that you are running a supported version of Cisco Integrated Management Controller (CIMC). The minimum that is supported and recommended versions of CIMC are listed in the "Compatibility" section of the [Release Notes](#) for your Cisco Nexus Dashboard release.

- Ensure that you have configured an IP address for the server's CIMC.

See [Configure a Cisco Integrated Management Controller IP address, on page 64](#).

- Ensure that Serial over LAN (SoL) is enabled in CIMC.

See [Enable Serial over LAN in the Cisco Integrated Management Controller, on page 65](#).

You might have a misconfiguration of SoL if the bootstrap fails at the `bootstrap peer nodes` point with this error:

```
Waiting for firstboot prompt on NodeX
```

- Ensure that all nodes are running the same release version image.
- If your Cisco Nexus Dashboard hardware came with a different release image than the one you want to deploy, we recommend deploying the cluster with the existing image first and then upgrading it to the needed release.

For example, if the hardware you received came with the release 3.2.1 image pre-installed, but you want to deploy release 4.3.1 instead, we recommend:

1. First, bring up the release 3.2.1 cluster, as described in the deployment guide for [that release](#).
2. Then upgrade to release 4.3.1, as described in [Upgrading a Nexus Dashboard 3.2.2 Cluster to This Release, on page 221](#).



Note

For brand new deployments, you can also choose to simply re-image the nodes with the latest version of the Cisco Nexus Dashboard (for example, if the hardware came with an image which does not support a direct upgrade to this release through the GUI workflow) before returning to this document for deploying the cluster. This process is described in the "Re-Imaging Nodes" section of the [Troubleshooting](#) article for this release.

- You must have at least a 1-node cluster.
- Extra secondary nodes can be added for horizontal scaling if required. For the maximum number of `secondary` and `standby` nodes in a single cluster, see the [Release Notes](#) for your release.

Configure a Cisco Integrated Management Controller IP address

Follow these steps to configure a Cisco Integrated Management Controller (CIMC) IP address.

Procedure

-
- Step 1** Power on the server.
- After the hardware diagnostic is complete, you will be prompted with different options controlled by the function (Fn) keys.
- Step 2** Press the **F8** key to enter the **Cisco IMC configuration Utility**.

Step 3 Follow these substeps.

- a) Set **NIC mode** to `Dedicated`.
- b) Choose between the **IPv4** and **IPv6** IP modes.

You can choose to enable or disable DHCP. If you disable DHCP, provide the static IP address, subnet, and gateway information.

- c) Ensure that **NIC Redundancy** is set to `None`.
- d) Press **F1** for more options such as hostname, DNS, default user passwords, port properties, and reset port profiles.

Step 4 Press **F10** to save the configuration and then restart the server.

Enable Serial over LAN in the Cisco Integrated Management Controller

Serial over LAN (SoL) is required for the `connect host` command, which you use to connect to a physical appliance node to provide basic configuration information. To use the SoL, you must first enable it on your Cisco Integrated Management Controller (CIMC).

Follow these steps to enable Serial over LAN in the Cisco Integrated Management Controller.

Procedure

Step 1 SSH into the node using the CIMC IP address and enter the sign-in credentials.

Step 2 Run these commands:

```
Server# scope sol
Server /sol # set enabled yes
Server /sol *# set baud-rate 115200
Server /sol *# commit
Server /sol *#
Server /sol # show
```

```
C220-WZP23150D4C# scope sol
C220-WZP23150D4C /sol # show
```

Enabled	Baud Rate(bps)	Com Port	SOL SSH Port
yes	115200	com0	2400

Step 3 In the command output, verify that `com0` is the com port for SoL.

This enables the system to monitor the console using the `connect host` command from the CIMC CLI, which is necessary for the cluster bringup.

Understanding and cabling physical appliances

These sections provide overview and cabling information for these supported Nexus Dashboard physical appliances.

ND-NODE-G5L (UCS-C225-M8)

These sections provide overview and cabling information for the ND-NODE-G5L (UCS-C225-M8) physical appliance.

Understanding the ND-NODE-G5L (UCS-C225-M8)

This section provides overview information for the ND-NODE-G5L (UCS-C225-M8) physical appliance.

The ND-NODE-G5L physical appliance is a Nexus Dashboard appliance that is based on the UCS C225 M8 server. Even though the Nexus Dashboard ND-NODE-G5L is based on the UCS C225 M8 server, you cannot replace components within the Nexus Dashboard ND-NODE-G5L as you would with the base UCS C225 M8 server; instead, you will perform a Return Material Authorization (RMA) operation for the appliance if a component within the Nexus Dashboard ND-NODE-G5L is faulty.

Because the Nexus Dashboard ND-NODE-G5L is based on the UCS C225 M8 server, you can refer to the [Cisco UCS C225 M8 Server Installation and Service Guide](#) for important UCS C225 M8 server information that is also applicable for the Nexus Dashboard ND-NODE-G5L, such as the information provided in these chapters in that document:

- Overview
- Installing the Server
- Server Specifications
- Storage Controller Considerations

However, because you cannot replace components within the Nexus Dashboard ND-NODE-G5L as you would with the base UCS C225 M8 server, these chapters from that document are not applicable for the Nexus Dashboard ND-NODE-G5L and should be ignored:

- Maintaining the Server
- Recycling Server Components
- Installation For Cisco UCS Manager Integration
- GPU Installation

The following sections provide information that is applicable for the Nexus Dashboard ND-NODE-G5L.

Overview

Cisco Nexus Dashboard provides a common platform for deploying Cisco Data Center applications. These applications provide real time analytics, visibility and assurance for policy and infrastructure.

The Cisco Nexus Dashboard server is required for installing and hosting the Cisco Nexus Dashboard application.

The appliance is orderable in the following versions:

- **ND-NODE-G5L**: Single-node appliance
- **ND-CLUSTERG5L**: Three-node version that leverages the same configuration as ND-NODE-G5L but with three appliances included

Components

The ND-NODE-G5L appliance is configured with the following components:

- **ND-M8-MLB**: ND M8 RACK MLB
- **ND-C225-M8S**: UCS C225 M8 Rack w/oCPU mem drv 1U wSFF HDD/SSD backplane
- **CIMC-LATEST-D**: IMC SW (Recommended) latest release for C-Series Servers
- **CON-L1NCO-UCSC5M8S**: CX LEVEL 1 8X7XNCDOS UCS C225 M8 Rack w/o CPU mem drv 1U w
- **ND-RAIL-D**: Ball Bearing Rail Kit for C220 & C240 M7/M8 rack servers
- **ND-CPU-A9655P**: ND AMD 9655P 2.6GHz 400W 96C/384MB Cache DDR5 6000MT/s
- **ND-M2-I240GB-D**: ND 240GB M.2 Boot SATA Intel SSD
- **ND-M2-HWRAID-D**: ND Cisco Boot optimized M.2 Raid controller
- **ND-TPM2-002D-D**: ND TPM 2.0 FIPS 140-2 MSW2022 compliant AMD M8 servers
- **ND-RIS1A-225M8**: ND C225 M8 1U Riser 1A PCIe Gen4 x16 HH
- **ND-HD24TB10KJ4-D**: ND 2.4TB 12G SAS 10K RPM SFF HDD (4Kn)
- **ND-SD960GBM3XEPD**: ND 960GB 2.5in Enter Perf 6G SATA Micron G2 SSD (3X)
- **ND-SD32TKA3XEP-D**: 3.2TB 2.5in Enter Perf 24G SAS Kioxia PM7 SSD (3X)
- **ND-M-V5Q50GV2-D**: Cisco VIC 15427 4x 10/25/50G mLOM C-Series w/Secure Boot
- **ND-FBRS2-C225M8**: C225 M8 Riser2 HH Filler Blank
- **ND-FBRS-C220-D**: C220M7 HH Riser3 blank
- Power supplies:
 - 1200W AC Titanium Power Supply for C-series Rack Servers
 - 1050W -48V DC Power Supply for UCS Rack Server
 - 1050W -48V DC Power Supply for APIC servers (India)
- **ND-MRX64G2RE5**: ND 64GB DDR5-6400 RDIMM 2Rx4 (16Gb)
- **ND-RAID-M1L16**: ND 24G Tri-Mode M1 RAID Controller w/4GB FBWC 16Drv
- **ND-P-ID10GC-D**: Cisco-Intel X710T2LG 2x10 GbE RJ45 PCIe NIC
- **ND-HSLP-C225M8**: ND C225 M8 Heatsink
- **ND-BBLKD-M8**: ND C-Series M6 & M8 SFF drive blanking panel
- **ND-DDR5-BLK**: ND DDR5 DIMM Blanks
- **ND-PSU-BLK-D**: Power Supply Blanking Panel for M7 / M8 servers
- **ND-HPBKT-225M8**: C225 M8 Tri-Mode 24G SAS RAID Controller Bracket
- **CBL-SAS-C225M8**: C225M8 SAS Cable Mainboard to RAID card

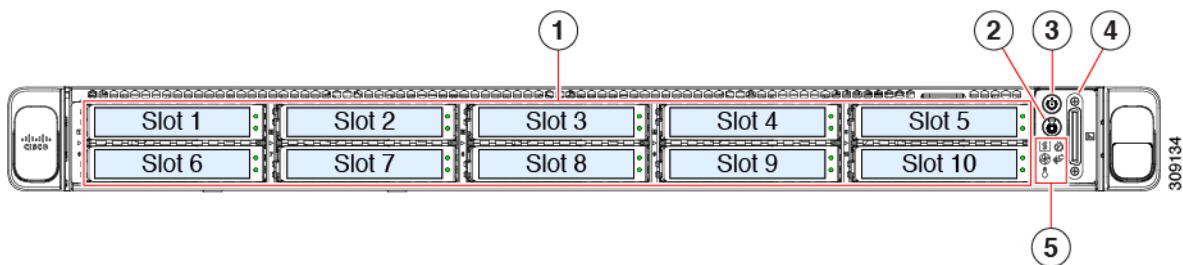
- **NO-POWER-CORD**: ECO friendly green option no power cable will be shipped
- **ND-DLOM-01-D**: Dedicated Mode BIOS setting for C-Series Servers
- **CNDL-DESELECT-D**: Conditional Deselect
- **OPTOUT-ENTL-SWAP**: License not needed: Entitlements updated in Smart Account

External Features

ND-NODE-G5L Front Panel Features

The following figure shows the front panel features of the small form-factor drive versions of the server.

Figure 6: ND-NODE-G5L Front Panel



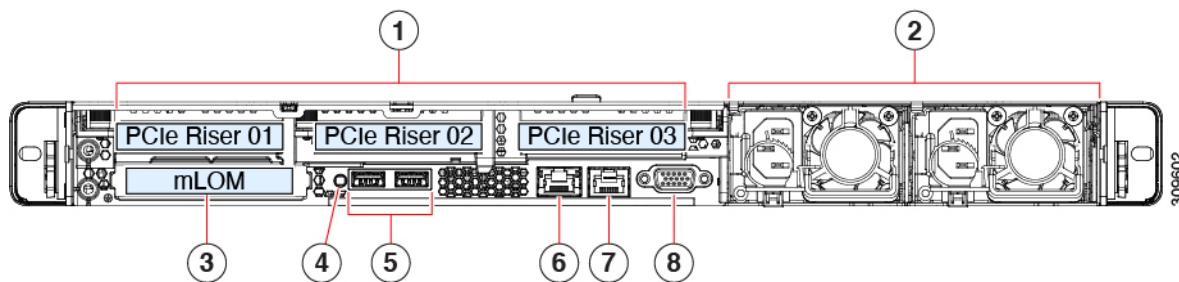
1	Drive bays 1 – 10 support SAS/SATA hard disk drives (HDDs) and solid state drives (SSDs).	2	Unit identification button/LED
3	Power button/power status LED	4	KVM connector (used with KVM cable that provides one DB-15 VGA, one DB-9 serial, and two USB 2.0 connectors)
5	System LED cluster: • Fan status LED • System Status LED • Power supply status LED • Network link activity LED • Temperature status LED		-

ND-NODE-G5L Rear Panel Features

The rear panel features can be different depending on the number and type of PCIe cards in the server.

The following figure shows the rear panel features of the server with three riser configuration.

Figure 7: ND-NODE-G5L Rear Panel Three-Riser Configuration



1	PCIe slots Following PCIe Riser combinations are available for 3 HH Riser cage configuration: • Riser 1: • Riser 1A (PCIe Gen4): Half-height, 3/4 length, x16.		
2	Power supply units (PSUs), two, which can be redundant when configured in 1+1 power mode.	3	Modular LAN-on-motherboard (mLOM) card bay (x16 PCIe lane)
4	System identification button/LED	5	USB 3.0 ports (two)
6	Dedicated 1 GB Ethernet management port	7	COM port (RJ-45 connector)
8	VGA video port (DB-15 connector)		

Status LEDs and Buttons

This section contains information for interpreting front, rear, and internal LED states.

Front-Panel LEDs

Figure 8: Front Panel LEDs

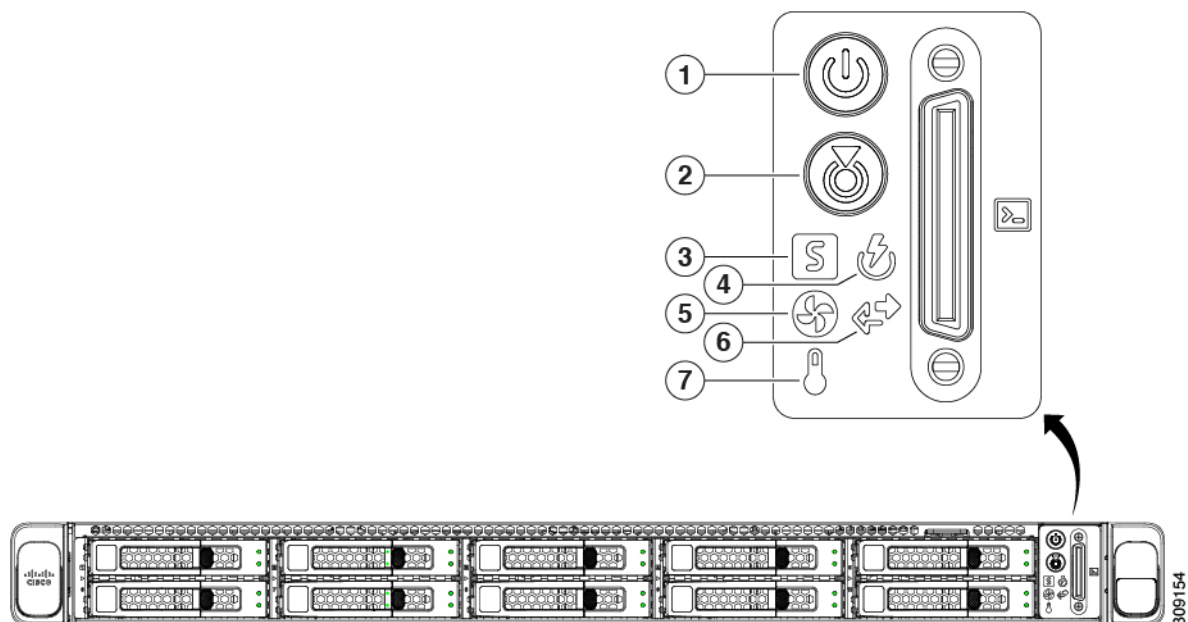


Table 13: Front Panel LEDs, Definition of States

	LED Name	States
1	Power button/LED ()	• Off—There is no AC power to the server. • Amber—The server is in standby power mode. Power is supplied only to the Cisco IMC and some motherboard functions. • Green—The server is in main power mode. Power is supplied to all server components.
2	Unit identification ()	• Off—The unit identification function is not in use. • Blue, blinking—The unit identification function is activated.

3	System health ()	• Green—The server is running in normal operating condition. • Green, blinking—The server is performing system initialization and memory check. • Amber, steady—The server is in a degraded operational state (minor fault). For example: • Power supply redundancy is lost. • CPUs are mismatched. • At least one CPU is faulty. • At least one DIMM is faulty. • At least one drive in a RAID configuration failed. • Amber, 2 blinks—There is a major fault with the system board. • Amber, 3 blinks—There is a major fault with the memory DIMMs. • Amber, 4 blinks—There is a major fault with the CPUs.
4	Power supply status ()	• Green—All power supplies are operating normally. • Amber, steady—One or more power supplies are in a degraded operational state. • Amber, blinking—One or more power supplies are in a critical fault state.
5	Fan status ()	• Green—All fan modules are operating properly. • Amber, blinking—One or more fan modules breached the non-recoverable threshold.
6	Network link activity ()	• Off—The Ethernet LOM port link is idle. • Green—One or more Ethernet LOM ports are link-active, but there is no activity. • Green, blinking—One or more Ethernet LOM ports are link-active, with activity.
7	Temperature status ()	• Green—The server is operating at normal temperature. • Amber, steady—One or more temperature sensors breached the critical threshold. • Amber, blinking—One or more temperature sensors breached the non-recoverable threshold.

Rear-Panel LEDs

Figure 9: Rear Panel LEDs

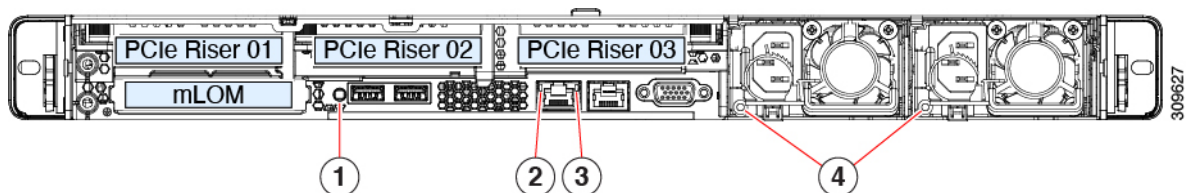


Table 14: Rear Panel LEDs, Definition of States

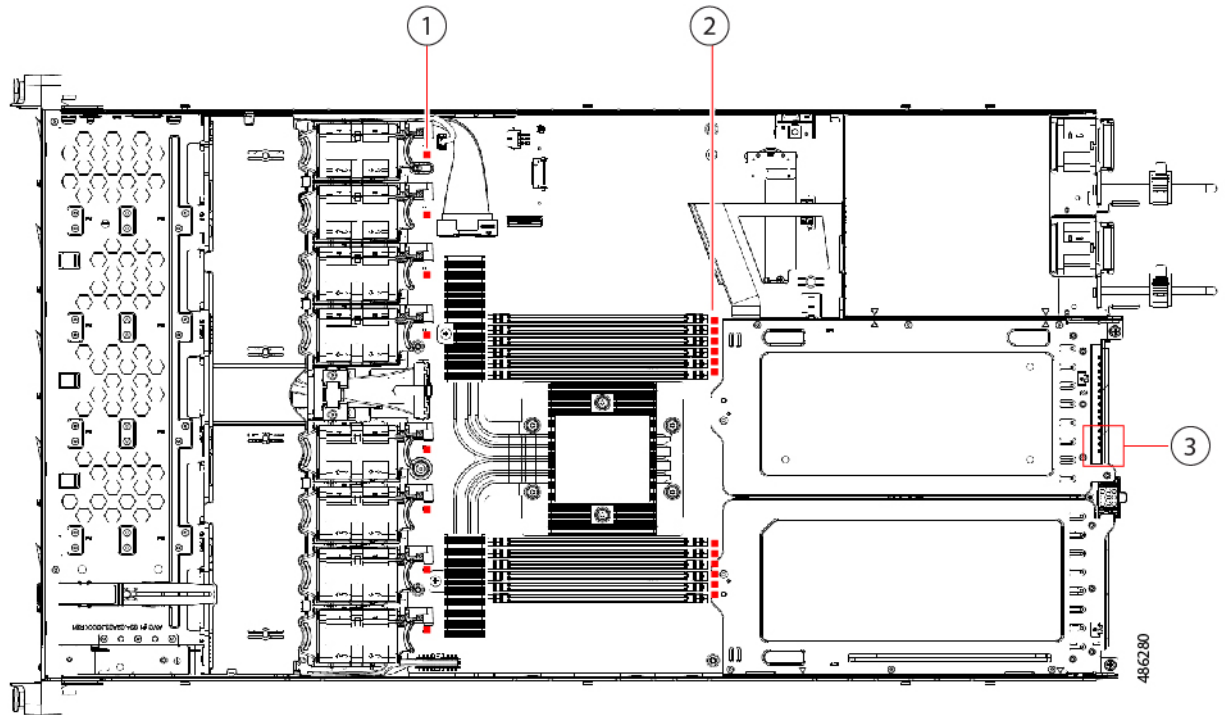
	LED Name	States
1	Rear unit identification	• Off—The unit identification function is not in use. • Blue, blinking—The unit identification function is activated.
2	1-Gb Ethernet dedicated management link speed	• Off—Link speed is 10 Mbps. • Amber—Link speed is 100 Mbps. • Green—Link speed is 1 Gbps.
3	1-Gb Ethernet dedicated management link status	• Off—No link is present. • Green—Link is active. • Green, blinking—Traffic is present on the active link.

4	Power supply status (one LED each power supply unit)	AC power supplies: • Off—No AC input (12 V main power off, 12 V standby power off). • Green, blinking—12 V main power off; 12 V standby power on. • Green, solid—12 V main power on; 12 V standby power on. • Amber, blinking—Warning threshold detected but 12 V main power on. • Amber, solid—Critical error detected; 12 V main power off (for example, over-current, over-voltage, or over-temperature failure). DC power supplies: • Off—No DC input (12 V main power off, 12 V standby power off). • Green, blinking—12 V main power off; 12 V standby power on. • Green, solid—12 V main power on; 12 V standby power on. • Amber, blinking—Warning threshold detected but 12 V main power on. • Amber, solid—Critical error detected; 12 V main power off (for example, over-current, over-voltage, or over-temperature failure).
---	--	---

Internal Diagnostic LEDs

The server has internal fault LEDs for CPUs, DIMMs, and fan modules.

Figure 10: Internal Diagnostic LED Locations

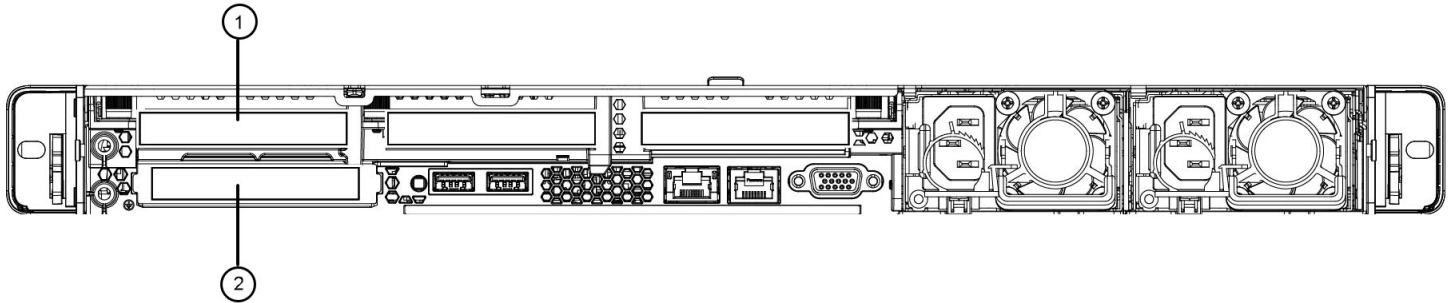


1	Fan module fault LEDs (one behind each fan connector on the motherboard) • Amber—Fan has a fault or is not fully seated. • Green—Fan is OK.	2	DIMM fault LEDs (one behind each DIMM socket on the motherboard) These LEDs operate only when the server is in standby power mode. • Amber—DIMM has a fault. • Off—DIMM is OK.
3	CPU fault LEDs (beside rear USB2 connector). These LEDs operate only when the server is in standby power mode. • Amber—CPU has a fault. • Off—CPU is OK.	-	

Physical node cabling

This figure shows the ND-NODE-G5L (UCS-C225-M8) physical server, where you will make these connections.

Figure 11: mLOM and PCIe riser 01 card used for node connectivity: ND-NODE-G5L (UCS-C225-M8)



1	Management connections: Through the two MGMT ports in the UCSC-O-ID10GC PCIe card installed in the PCIe Riser 01 location.
2	Data connections: Ports are numbered 1, 2, 3, and 4, from left-to-right in the UCSC-M-V5Q50GV2-D (Cisco UCS VIC 15427 Quad Port 10/25/50G CNA) Modular LAN-on-motherboard (mLOM). See the "Data network connections" information below for the supported port channel configurations.



Note The OCP card included on the ND-NODE-G5L servers supports a 1Gb copper connection for management only. All other network connections for Nexus Dashboard need to leverage the four port VIC card (callout 2 in the figure above). This VIC card supports 10/25/50Gbps, and the recommended SFP+ cable is the SFP-10G-AOC3M, but Cisco also offers 5-meter and 7-meter options as well. The VIC card requires a minimum of two connections per server for data network connectivity. These VIC connections can leverage any supported SFP, but Cisco recommends this connection for seamless deployments of Nexus Dashboard.

The physical nodes can be deployed with these guidelines:

- The ND-NODE-G5L comes with a UCSC-O-ID10GC PCIe card (shown in the above diagram), which you use for Nexus Dashboard management network connectivity.
- The ND-NODE-G5L also comes with a UCSC-M-V5Q50GV2-D (Cisco UCS VIC 15427 Quad Port 10/25/50G CNA) Modular LAN on Motherboard (mLOM) card, which you use to connect to the Nexus Dashboard data network.



Note Verify the Cisco UCS VIC card port mapping for the data interfaces using the information provided in "Data network connections" before cabling the node. The data interfaces must connect to individual host-facing switch ports; do not configure switch port channels or vPC for these connections.

- Verify that the port-channel feature is enabled on the VIC.
 - For the ND-NODE-L4 (UCS-C225-M6) and ND-NODE-G5S/ND-NODE-G5L (UCS-C225-M8), the port-channel feature is enabled by default at the factory. No additional configurations are necessary for those platforms, but you can use these procedures to verify that the port-channel feature is enabled correctly and enable it on those platforms, if necessary.

- For the `SE-NODE-G2` (UCS-C220-M5), you must manually enable the port-channel feature on the VIC.

1. Log into the CIMC.
2. Click the hamburger icon, then navigate to:
Networking > Adapter Card 1 > General
3. In the **Adapter Card Properties** area, locate the **Port Channel** field.

If the **Port Channel** option is disabled (has no check in the check box), enable the feature (put a check in the check box) before beginning the Nexus Dashboard deployment process.

When connecting the node to your management and data networks:

- The interfaces are configured as Linux bonds, one for the data interfaces (bond0) and one for the management interfaces (bond1), running in Active-Standby mode.
- **Management network connections:**
 - You must use the `mgmt0` and `mgmt1` on the PCIe card.
 - All ports must have the same speed, either 1G or 10G.
- **Data network connections:**
 - On the `ND-NODE-G5L` server, you must use optical connections through the necessary port channel combinations in the UCSC-M-V5Q50GV2-D (Cisco UCS VIC 15427 Quad Port 10/25/50G CNA) PCIe card.



Note For 25/50 GB speed connections, you will need one of the following pairs of Forward Error Correction (FEC) configurations:

On the N9000	On the Nexus Dashboard appliance (via CIMC)
FEC AUTO	c174
FC-FEC	c174
FEC OFF	FEC OFF

- All ports must have the same speed, either 10G, 25G, or 50G.
- Attach the cables for the data network connections, where one cable connects to `fabric0` and the second cable connects to `fabric1`.

`fabric0` and `fabric1` in the `ND-NODE-G5L` server corresponds to these ports:

- Port-1 and Port-2 correspond to `fabric0`
- Port-3 and Port-4 correspond to `fabric1`

You can therefore have any one of these port channel combinations:

First data cable (<code>fabric0</code>)	Second data cable (<code>fabric1</code>)
Port-1	Port-3
Port-1	Port-4
Port-2	Port-3
Port-2	Port-4

See [Figure 11: mLOM and PCIe riser 01 card used for node connectivity: ND-NODE-G5L \(UCS-C225-M8\)](#), on page 75 for the port locations.

You can use both `fabric0` and `fabric1` for data network connectivity as Active-Standby.



Caution

- Do not configure a switch port channel or vPC for these connections. The supported pairings listed above identify the valid physical port mappings on the VIC card only.
 - If you connect the two cables for the data network connections using different port channel combinations from those listed above, there will be MAC move notifications on the upstream switch and the ports will flap.
-
- If you connect the nodes to Cisco Catalyst switches, packets are tagged on those Catalyst switches with `vlan0` if no VLAN is specified. In this case, you must add `switchport voice vlan dot1p` command to the switch interfaces where the nodes are connected to ensure reachability over the data network.

What's next

Go to [Deploy Nexus Dashboard as a physical appliance, on page 110](#) to deploy the Nexus Dashboard.

ND-NODE-G5S (UCS-C225-M8)

These sections provide overview and cabling information for the ND-NODE-G5S (UCS-C225-M8) physical appliance.

Understanding the ND-NODE-G5S (UCS-C225-M8)

This section provides overview information for the ND-NODE-G5S (UCS-C225-M8) physical appliance.

The ND-NODE-G5S physical appliance is a Nexus Dashboard appliance that is based on the UCS C225 M8 server. Even though the Nexus Dashboard ND-NODE-G5S is based on the UCS C225 M8 server, you cannot replace components within the Nexus Dashboard ND-NODE-G5S as you would with the base UCS C225 M8 server; instead, you will perform a Return Material Authorization (RMA) operation for the appliance if a component within the Nexus Dashboard ND-NODE-G5S is faulty.

Because the Nexus Dashboard ND-NODE-G5S is based on the UCS C225 M8 server, you can refer to the [Cisco UCS C225 M8 Server Installation and Service Guide](#) for important UCS C225 M8 server information that is

also applicable for the Nexus Dashboard `ND-NODE-G5S`, such as the information provided in these chapters in that document:

- Overview
- Installing the Server
- Server Specifications
- Storage Controller Considerations

However, because you cannot replace components within the Nexus Dashboard `ND-NODE-G5S` as you would with the base UCS C225 M8 server, these chapters from that document are not applicable for the Nexus Dashboard `ND-NODE-G5S` and should be ignored:

- Maintaining the Server
- Recycling Server Components
- Installation For Cisco UCS Manager Integration
- GPU Installation

The following sections provide information that is applicable for the Nexus Dashboard `ND-NODE-G5S`.

Overview

Cisco Nexus Dashboard provides a common platform for deploying Cisco Data Center applications. These applications provide real time analytics, visibility and assurance for policy and infrastructure.

The Cisco Nexus Dashboard server is required for installing and hosting the Cisco Nexus Dashboard application.

The appliance is orderable in the following versions:

- **ND-NODE-G5S**: Single-node appliance
- **ND-CLUSTERG5S**: Three-node version that leverages the same configuration as ND-NODE-G5S but with three appliances included

Components

The ND-NODE-G5S appliance is configured with the following components:

- **CIMC-LATEST-D**: IMC SW (Recommended) latest release for C-Series Servers
- **ND-CPU-A9454P**: ND AMD 9454P 2.75GHz 290W 48C/256MB Cache DDR5 4800MT/s
- **ND-M2-240G-D**: ND 240GB M.2 SATA Micron G2 SSD
- **ND-M2-HWRAID-D**: ND Cisco Boot optimized M.2 Raid controller
- **ND-TPM2-002D-D**: ND TPM 2.0 FIPS 140-2 MSW2022 compliant AMD M8 servers
- **ND-RIS1A-225M8**: ND C225 M8 1U Riser 1A PCIe Gen4 x16 HH
- **ND-HD24TB10KJ4-D**: ND 2.4TB 12G SAS 10K RPM SFF HDD (4Kn)
- **ND-SD960GBM3XEPD**: ND 960GB 2.5in Enter Perf 6G SATA Micron G2 SSD (3X)
- **ND-RAIL-D**: Ball Bearing Rail Kit for C220 & C240 M7/M8 rack servers

- Power supplies:
 - 1200W AC Titanium Power Supply for C-series Rack Servers
 - 1050W -48V DC Power Supply for UCS Rack Server
 - 1050W -48V DC Power Supply for APIC servers (India)
- **ND-MRX32G1RE3**: ND 32GB DDR5-5600 RDIMM 1Rx4 (16Gb)
- **ND-RAID-M1L16**: ND 24G Tri-Mode M1 RAID Controller w/4GB FBWC 16Drv
- **ND-O-ID10GC-D**: Intel X710T2LOCPV3G1L 2x10GbE RJ45 OCP3.0 NIC
- **ND-OC3-KIT-D**: C2XX OCP 3.0 Interposer W/Mech Assy
- **ND-P-V5Q50G-D**: Cisco VIC 15425 4x 10/25/50G PCIe C-Series w/Secure Boot

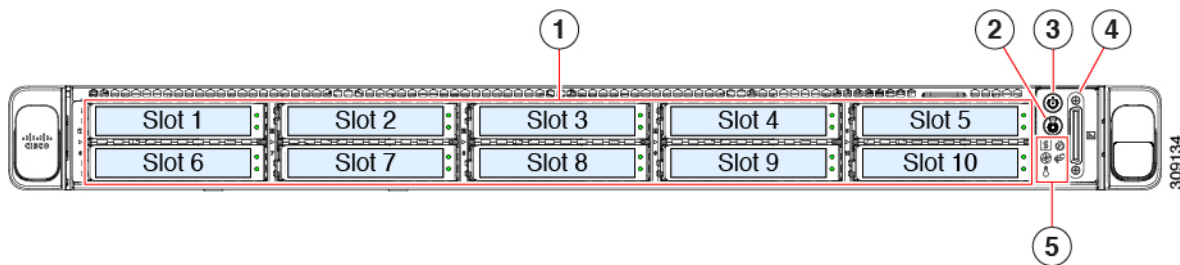
External Features

ND-NODE-G5S Front Panel Features

The following figure shows the front panel features of the small form-factor drive versions of the server.

For definitions of LED states, see [Front-Panel LEDs](#), on page 70.

Figure 12: ND-NODE-G5S Front Panel



1	Drive bays 1 – 10 support SAS/SATA hard disk drives (HDDs) and solid state drives (SSDs).	2	Unit identification button/LED
3	Power button/power status LED	4	KVM connector (used with KVM cable that provides one DB-15 VGA, one DB-9 serial, and two USB 2.0 connectors)
5	System LED cluster: • Fan status LED • System Status LED • Power supply status LED • Network link activity LED • Temperature status LED		-

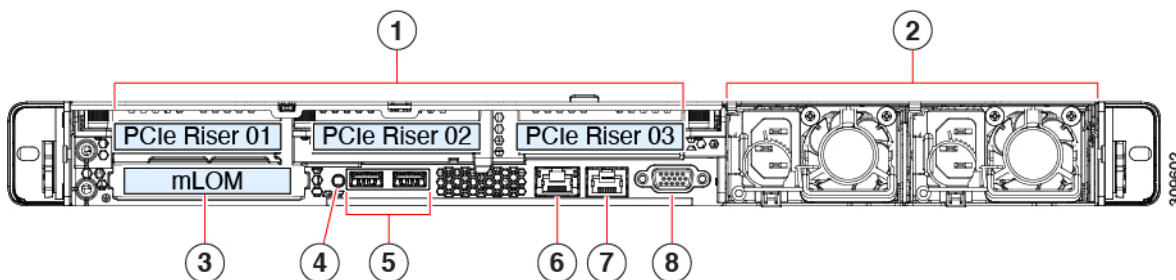
ND-NODE-G5S Rear Panel Features

The rear panel features can be different depending on the number and type of PCIe cards in the server.

The following figure shows the rear panel features of the server with three riser configuration.

For definitions of LED states, see [Rear-Panel LEDs, on page 72](#).

Figure 13: ND-NODE-G5S Rear Panel Three-Riser Configuration



1	PCIe slots Following PCIe Riser combinations are available for 3 HH Riser cage configuration: • Riser 1: • Riser 1A (PCIe Gen4): Half-height, 3/4 length, x16, NCSI, Single Wide GPU.		
2	Power supply units (PSUs), two, which can be redundant when configured in 1+1 power mode.	3	Modular LAN-on-motherboard (mLOM) card bay (x16 PCIe lane)
4	System identification button/LED	5	USB 3.0 ports (two)
6	Dedicated 1 GB Ethernet management port	7	COM port (RJ-45 connector)
8	VGA video port (DB-15 connector)		

PCIe Slot Specifications

The following tables describe the specifications for the slots in three riser combination.

Table 15: PCIe Riser 1

Slot Number	Electrical Lane Width	Connector Length	Maximum Card Length	Card Height (Rear Panel Opening)	NCSI Support
1A	Gen4 x16	x24 connector	¾ length	Half-height	Yes
1B	Gen5 x16	x24 connector	¾ length	Half-height	Yes

Status LEDs and Buttons

This section contains information for interpreting front, rear, and internal LED states.

Front-Panel LEDs

Figure 14: Front Panel LEDs

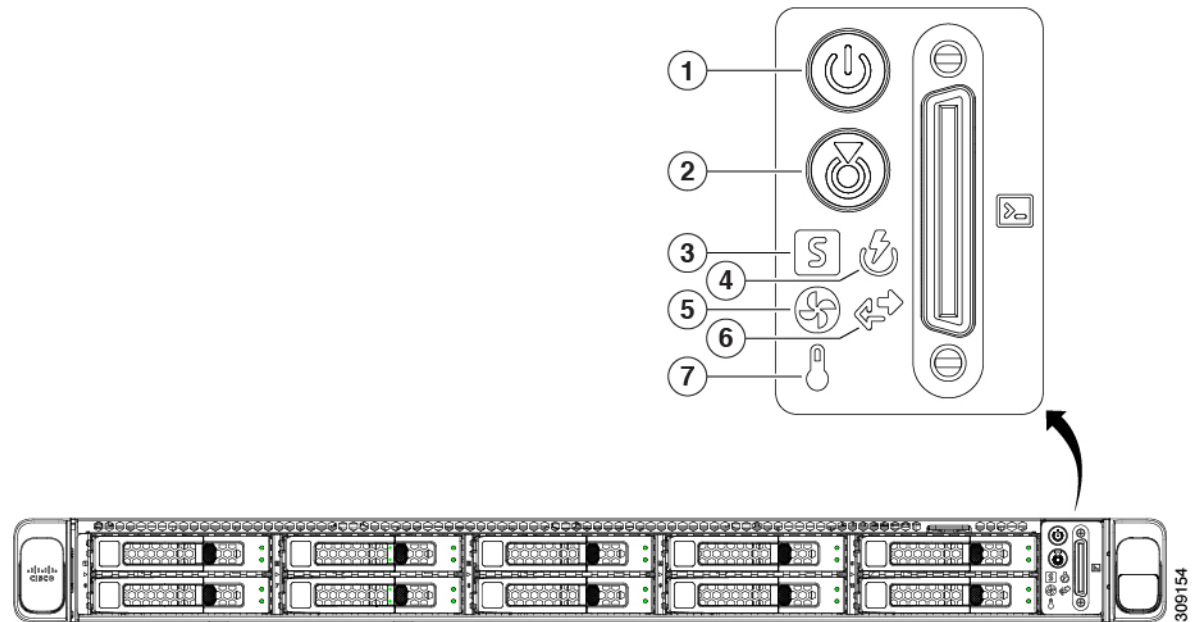


Table 16: Front Panel LEDs, Definition of States

	LED Name	States
1	Power button/LED ()	• Off—There is no AC power to the server. • Amber—The server is in standby power mode. Power is supplied only to the Cisco IMC and some motherboard functions. • Green—The server is in main power mode. Power is supplied to all server components.
2	Unit identification ()	• Off—The unit identification function is not in use. • Blue, blinking—The unit identification function is activated.

3	System health ()	• Green—The server is running in normal operating condition. • Green, blinking—The server is performing system initialization and memory check. • Amber, steady—The server is in a degraded operational state (minor fault). For example: • Power supply redundancy is lost. • CPUs are mismatched. • At least one CPU is faulty. • At least one DIMM is faulty. • At least one drive in a RAID configuration failed. • Amber, 2 blinks—There is a major fault with the system board. • Amber, 3 blinks—There is a major fault with the memory DIMMs. • Amber, 4 blinks—There is a major fault with the CPUs.
4	Power supply status ()	• Green—All power supplies are operating normally. • Amber, steady—One or more power supplies are in a degraded operational state. • Amber, blinking—One or more power supplies are in a critical fault state.
5	Fan status ()	• Green—All fan modules are operating properly. • Amber, blinking—One or more fan modules breached the non-recoverable threshold.
6	Network link activity ()	• Off—The Ethernet LOM port link is idle. • Green—One or more Ethernet LOM ports are link-active, but there is no activity. • Green, blinking—One or more Ethernet LOM ports are link-active, with activity.
7	Temperature status ()	• Green—The server is operating at normal temperature. • Amber, steady—One or more temperature sensors breached the critical threshold. • Amber, blinking—One or more temperature sensors breached the non-recoverable threshold.

Rear-Panel LEDs

Figure 15: Rear Panel LEDs

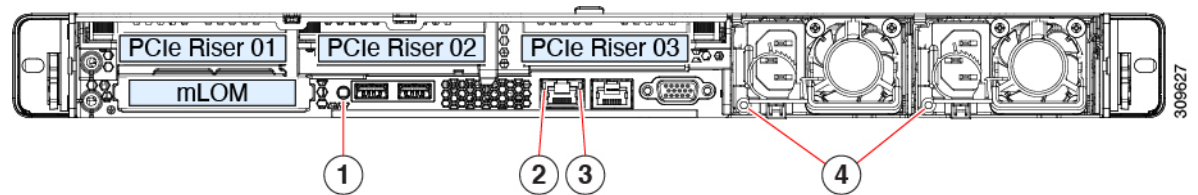


Table 17: Rear Panel LEDs, Definition of States

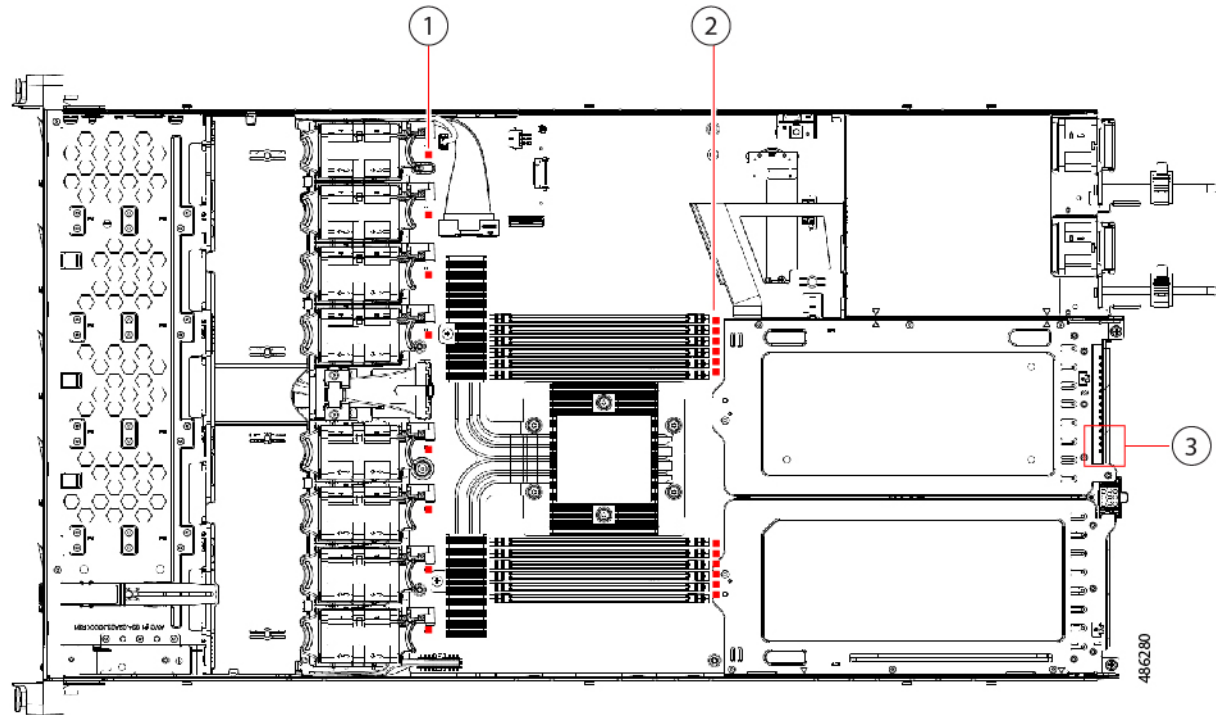
	LED Name	States
1	Rear unit identification	• Off—The unit identification function is not in use. • Blue, blinking—The unit identification function is activated.
2	1-Gb Ethernet dedicated management link speed	• Off—Link speed is 10 Mbps. • Amber—Link speed is 100 Mbps. • Green—Link speed is 1 Gbps.
3	1-Gb Ethernet dedicated management link status	• Off—No link is present. • Green—Link is active. • Green, blinking—Traffic is present on the active link.

4	Power supply status (one LED each power supply unit)	AC power supplies: • Off—No AC input (12 V main power off, 12 V standby power off). • Green, blinking—12 V main power off; 12 V standby power on. • Green, solid—12 V main power on; 12 V standby power on. • Amber, blinking—Warning threshold detected but 12 V main power on. • Amber, solid—Critical error detected; 12 V main power off (for example, over-current, over-voltage, or over-temperature failure). DC power supplies: • Off—No DC input (12 V main power off, 12 V standby power off). • Green, blinking—12 V main power off; 12 V standby power on. • Green, solid—12 V main power on; 12 V standby power on. • Amber, blinking—Warning threshold detected but 12 V main power on. • Amber, solid—Critical error detected; 12 V main power off (for example, over-current, over-voltage, or over-temperature failure).
---	--	---

Internal Diagnostic LEDs

The server has internal fault LEDs for CPUs, DIMMs, and fan modules.

Figure 16: Internal Diagnostic LED Locations

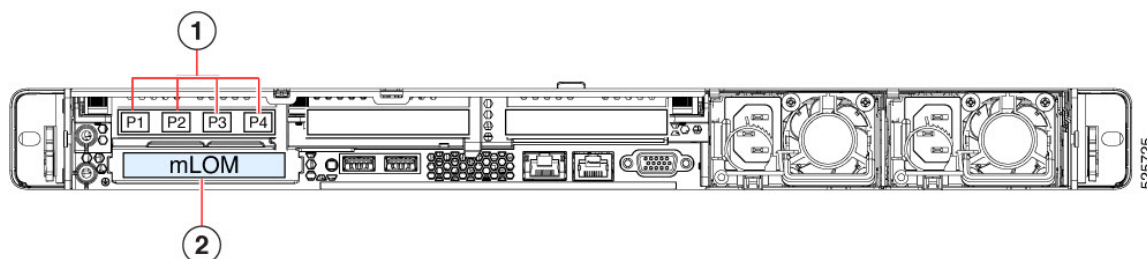


1	Fan module fault LEDs (one behind each fan connector on the motherboard) • Amber—Fan has a fault or is not fully seated. • Green—Fan is OK.	2	DIMM fault LEDs (one behind each DIMM socket on the motherboard) These LEDs operate only when the server is in standby power mode. • Amber—DIMM has a fault. • Off—DIMM is OK.
3	CPU fault LEDs (beside rear USB2 connector). These LEDs operate only when the server is in standby power mode. • Amber—CPU has a fault. • Off—CPU is OK.	-	

Physical node cabling

This figure shows the ND-NODE-G5S (UCS-C225-M8) physical server, where you will make these connections.

Figure 17: mLOM and PCIe riser 01 card used for node connectivity: ND-NODE-G5S (UCS-C225-M8)



1	Data connections: Ports are numbered 1, 2, 3, and 4, from left-to-right in the UCSC-P-V5Q50G-D (Cisco UCS VIC 15425 Quad Port 10/25/50G CNA) PCIE card installed in the PCIe Riser 01 location. See the "Data network connections" information below for the supported port channel configurations.
2	Management connections: Through the two MGMT ports in the Modular LAN-on-motherboard (mLOM).



Note The OCP card included on the ND-NODE-G5S servers supports a 1Gb copper connection for management only. All other network connections for Nexus Dashboard need to leverage the four port VIC card (callout 1 in the figure above). This VIC card supports 10/25/50Gbps, and the recommended SFP+ cable is the SFP-10G-AOC3M, but Cisco also offers 5-meter and 7-meter options as well. The VIC card requires a minimum of two connections per server for data network connectivity. These VIC connections can leverage any supported SFP, but Cisco recommends this connection for seamless deployments of Nexus Dashboard.

The physical nodes can be deployed with these guidelines:

- All servers come with a Modular LAN on Motherboard (mLOM) card, which you use to connect to the Nexus Dashboard management network.
- The ND-NODE-G5S includes a UCSC-P-V5Q50G-D (Cisco UCS VIC 15425 Quad Port 10/25/50G CNA) PCIE card in the "PCIe-Riser-01" slot (shown in the above diagram), which you use for Nexus Dashboard data network connectivity.



Note Verify the Cisco UCS VIC card port mapping for the data interfaces using the information provided in "Data network connections" before cabling the node. The data interfaces must connect to individual host-facing switch ports; do not configure switch port channels or vPC for these connections.

- Verify that the port-channel feature is enabled on the VIC.
 - For the ND-NODE-L4 (UCS-C225-M6) and ND-NODE-G5S/ND-NODE-G5L (UCS-C225-M8), the port-channel feature is enabled by default at the factory. No additional configurations are necessary for those platforms, but you can use these procedures to verify that the port-channel feature is enabled correctly and enable it on those platforms, if necessary.

- For the `SE-NODE-G2` (UCS-C220-M5), you must manually enable the port-channel feature on the VIC.

1. Log into the CIMC.
2. Click the hamburger icon, then navigate to:
Networking > Adapter Card 1 > General
3. In the **Adapter Card Properties** area, locate the **Port Channel** field.

If the **Port Channel** option is disabled (has no check in the check box), enable the feature (put a check in the check box) before beginning the Nexus Dashboard deployment process.

When connecting the node to your management and data networks:

- The interfaces are configured as Linux bonds, one for the data interfaces (bond0) and one for the management interfaces (bond1), running in Active-Standby mode.
- **Management network connections:**
 - You must use the `mgmt0` and `mgmt1` on the mLOM card.
 - All ports must have the same speed, either 1G or 10G.
- **Data network connections:**
 - On the `ND-NODE-G5S` server, you must use optical connections through the necessary port channel combinations in the UCSC-P-V5Q50G-D (Cisco UCS VIC 15425 Quad Port 10/25/50G CNA) PCIE card.



Note For 25/50 GB speed connections, you will need one of the following pairs of Forward Error Correction (FEC) configurations:

On the N9000	On the Nexus Dashboard appliance (via CIMC)
FEC AUTO	cl74
FC-FEC	cl74
FEC OFF	FEC OFF

- All ports must have the same speed, either 10G, 25G, or 50G.
- Attach the cables for the data network connections, where one cable connects to `fabric0` and the second cable connects to `fabric1`.

`fabric0` and `fabric1` in the `ND-NODE-G5S` server corresponds to these ports:

- Port-1 and Port-2 correspond to `fabric0`
- Port-3 and Port-4 correspond to `fabric1`

You can therefore have any one of these port channel combinations:

First data cable (<i>fabric0</i>)	Second data cable (<i>fabric1</i>)
Port-1	Port-3
Port-1	Port-4
Port-2	Port-3
Port-2	Port-4

See [Figure 17: mLOM and PCIe riser 01 card used for node connectivity: ND-NODE-G5S \(UCS-C225-M8\)](#), on page 86 for the port locations.

You can use both *fabric0* and *fabric1* for data network connectivity as Active-Standby.



Caution

- Do not configure a switch port channel or vPC for these connections. The supported pairings listed above identify the valid physical port mappings on the VIC card only.
 - If you connect the two cables for the data network connections using different port channel combinations from those listed above, there will be MAC move notifications on the upstream switch and the ports will flap.
-
- If you connect the nodes to Cisco Catalyst switches, packets are tagged on those Catalyst switches with *vlan0* if no VLAN is specified. In this case, you must add `switchport voice vlan dot1p` command to the switch interfaces where the nodes are connected to ensure reachability over the data network.

What's next

Go to [Deploy Nexus Dashboard as a physical appliance, on page 110](#) to deploy the Nexus Dashboard.

ND-NODE-L4 (UCS-C225-M6)

These sections provide overview and cabling information for the ND-NODE-L4 (UCS-C225-M6) physical appliance.

Understanding the ND-NODE-L4 (UCS-C225-M6)

This section provides overview information for the ND-NODE-L4 (UCS-C225-M6) physical appliance.

The ND-NODE-L4 physical appliance is a Nexus Dashboard appliance that is based on the UCS-C225-M6 server. Even though the Nexus Dashboard ND-NODE-L4 is based on the UCS-C225-M6 server, you cannot replace components within the Nexus Dashboard ND-NODE-L4 as you would with the base UCS-C225-M6 server; instead, you will perform a Return Material Authorization (RMA) operation for the appliance if a component within the Nexus Dashboard ND-NODE-L4 is faulty.

Because the Nexus Dashboard ND-NODE-L4 is based on the UCS-C225-M6 server, you can refer to the [Cisco UCS C225 M6 Server Installation and Service Guide](#) for important UCS-C225-M6 server information that

is also applicable for the Nexus Dashboard ND-NODE-L4, such as the information provided in these chapters in that document:

- Overview
- Installing the Server
- Server Specifications
- Storage Controller Considerations

However, because you cannot replace components within the Nexus Dashboard ND-NODE-L4 as you would with the base UCS-C225-M6 server, these chapters from that document are not applicable for the Nexus Dashboard ND-NODE-L4 and should be ignored:

- Maintaining the Server
- Installation For Cisco UCS Manager Integration
- GPU Installation

The following sections provide information that is applicable for the Nexus Dashboard ND-NODE-L4.

Overview

Cisco Nexus Dashboard provides a common platform for deploying Cisco Data Center applications. These applications provide real time analytics, visibility and assurance for policy and infrastructure.

The Cisco Nexus Dashboard server is required for installing and hosting the Cisco Nexus Dashboard application.

The server is orderable in the following version:

- ND-NODE-L4 — Small form-factor (SFF) drives, with 10-drive backplane. Supports up to 10 2.5-inch SAS/SATA drives. Drive bays 1 and 2 support NVMe SSDs.

Following PCIe Riser combinations are available:

- One half-height riser card in PCIe Riser 1
- Three half-height riser cards in PCIe Riser 1, 2, 3
- Two full-height riser cards Riser 1 and 3
- Riser 1—Supports Riser1. Supports single x16 PCIe supporting full height 3/4 length cards in 2 riser configuration (or) Half-height 3/4-length cards in 3 riser configuration and NC-SI from Pilot4.
- Riser 2—Supports Riser 1. Supports single x16 PCIe supporting only Half-height 3/4-length cards in 3-riser configuration.
- Riser 3—Supports Riser 3A, 3B. PCIe slot 3 with the following options:
 - Riser3A Supports single x16 PCIe supporting half height 3/4 length cards in 3 riser configuration and NC-SI.
 - Riser3B Supports single x16 PCIe supporting full height 3/4-length cards in 2 riser configuration and NC-SI.
- 2 10GBase-T Ethernet LAN over Motherboard (LOM) ports for network connectivity, plus one 1 Gigabit Ethernet dedicated management port

- One mLOM/VIC card provides 10G/25G/40G/50G/100G connectivity. Supported cards are:
 - Cisco VIC 1455 VIC PCIE – Quad Port 10/25G SFP28 (UCSC-PCIE-C25Q-04)

External Features

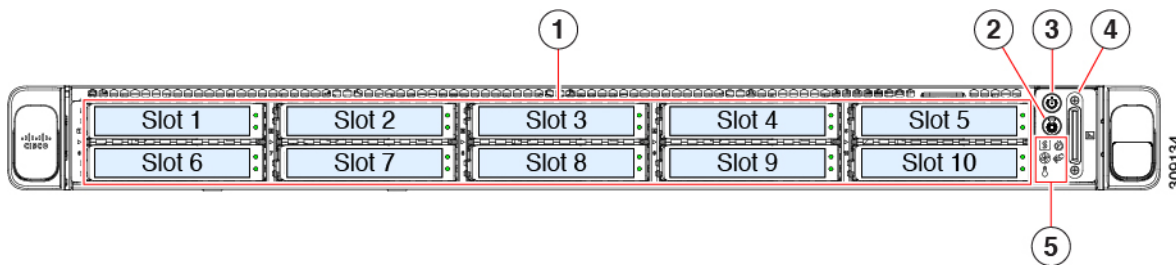
This topic shows the external features of the server versions.

Cisco ND-NODE-L4 (SFF Drives) Front Panel Features

The following figure shows the front panel features of the small form-factor drive versions of the server.

For definitions of LED states, see [Front-Panel LEDs](#), on page 92.

Figure 18: ND-NODE-L4 (SFF Drives) Front Panel



1 UCSC-C225-M6S Version—Drive bays 1 – 10 support SAS/SATA hard disk drives (HDDs) and solid state drives (SSDs). As an option, drive bays 1-4 can contain up to 4 NVMe drives in any number up to 4. Drive bays 5 through 10 support only SAS/SATA HDDs or SSDs. UCSC-C225-M6N Version—Drive bays 1—10 supports 2.5-inch NVMe-only SSDs.	2 Unit identification button/LED
3 Power button/power status LED	4 KVM connector (used with KVM cable that provides one DB-15 VGA, one DB-9 serial, and two USB 2.0 connectors)
5 System LED cluster: • Fan status LED • System Status LED • Power supply status LED • Network link activity LED • Temperature status LED	-

Cisco ND-NODE-L4 Rear Panel Features

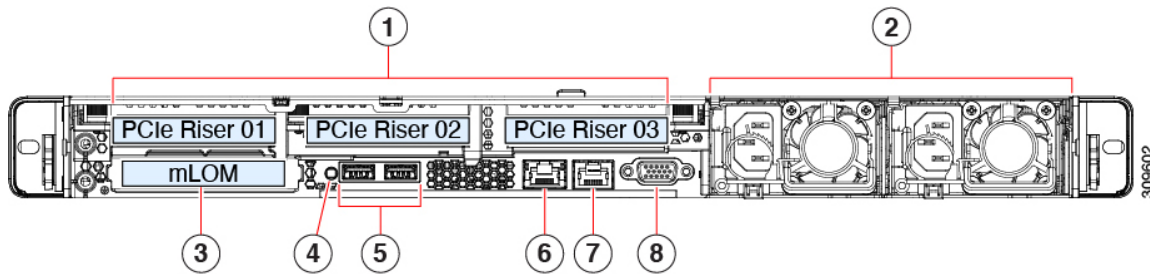
The rear panel features can be different depending on the number and type of PCIe cards in the server.

By default, single CPU servers come with only one half-height riser 1 installed, and dual CPU servers support all three half-height risers.

The following figure shows the rear panel features of the server with three riser configuration.

For definitions of LED states, see [Rear-Panel LEDs](#), on page 94.

Figure 19: Cisco ND-NODE-L4 Rear Panel Three Riser Configuration



The following figure shows the rear panel features of the server with three riser configuration.

PCIe slots Following PCIe Riser combinations are available: • One half-height riser card in PCIe Riser 1	
2 Power supply units (PSUs), two which can be redundant when configured in 1+1 power mode.	3 Modular LAN-on-motherboard (mLOM) card bay (x16 PCIe lane)
4 System identification button/LED	5 USB 3.0 ports (two)
6 Dedicated 1 GB Ethernet management port	7 LOM port (RJ-45 connector)
8 VGA video port (DB-15 connector)	

PCIe Slot Specifications

The following tables describe the specifications for the slots in three riser combination.

Table 18: PCIe Riser 1

Slot Number	Electrical Lane Width	Connector Length	Maximum Card Length	Card Height (Rear Panel Opening)	NCSI Support
1	Gen-3 and 4 x16	x24 connector	$\frac{3}{4}$ length	Half-height	Yes

Status LEDs and Buttons

This section contains information for interpreting front, rear, and internal LED states.

Front-Panel LEDs

Figure 20: Front Panel LEDs

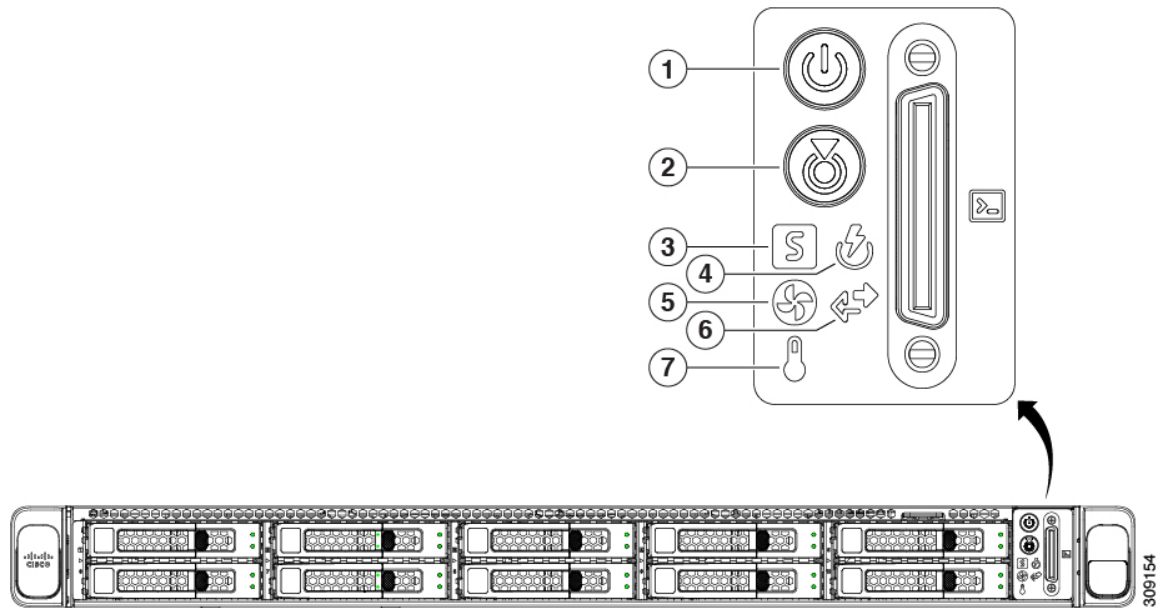


Table 19: Front Panel LEDs, Definition of States

LED Name	States
1 Power button/LED ()	• Off—There is no AC power to the server. • Amber—The server is in standby power mode. Power is supplied only to the Cisco IMC and some motherboard functions. • Green—The server is in main power mode. Power is supplied to all server components.
2 Unit identification ()	• Off—The unit identification function is not in use. • Blue, blinking—The unit identification function is activated.

3 System health ()	• Green—The server is running in normal operating condition. • Green, blinking—The server is performing system initialization and memory check. • Amber, steady—The server is in a degraded operational state (minor fault). For example: • Power supply redundancy is lost. • CPUs are mismatched. • At least one CPU is faulty. • At least one DIMM is faulty. • At least one drive in a RAID configuration failed. • Amber, 2 blinks—There is a major fault with the system board. • Amber, 3 blinks—There is a major fault with the memory DIMMs. • Amber, 4 blinks—There is a major fault with the CPUs.
4 Power supply status ()	• Green—All power supplies are operating normally. • Amber, steady—One or more power supplies are in a degraded operational state. • Amber, blinking—One or more power supplies are in a critical fault state.
5 Fan status ()	• Green—All fan modules are operating properly. • Amber, blinking—One or more fan modules breached the non-recoverable threshold.
6 Network link activity ()	• Off—The Ethernet LOM port link is idle. • Green—One or more Ethernet LOM ports are link-active, but there is no activity. • Green, blinking—One or more Ethernet LOM ports are link-active, with activity.
7 Temperature status ()	• Green—The server is operating at normal temperature. • Amber, steady—One or more temperature sensors breached the critical threshold. • Amber, blinking—One or more temperature sensors breached the non-recoverable threshold.

Rear-Panel LEDs

Figure 21: Rear Panel LEDs

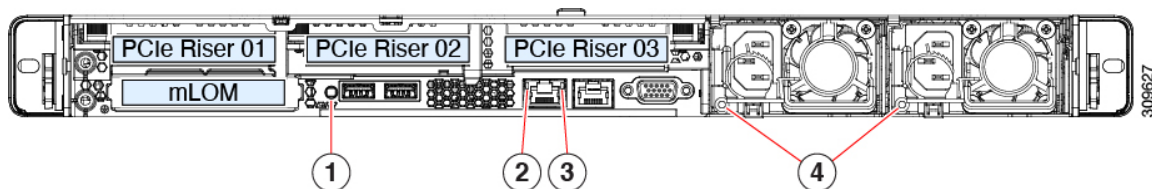


Table 20: Rear Panel LEDs, Definition of States

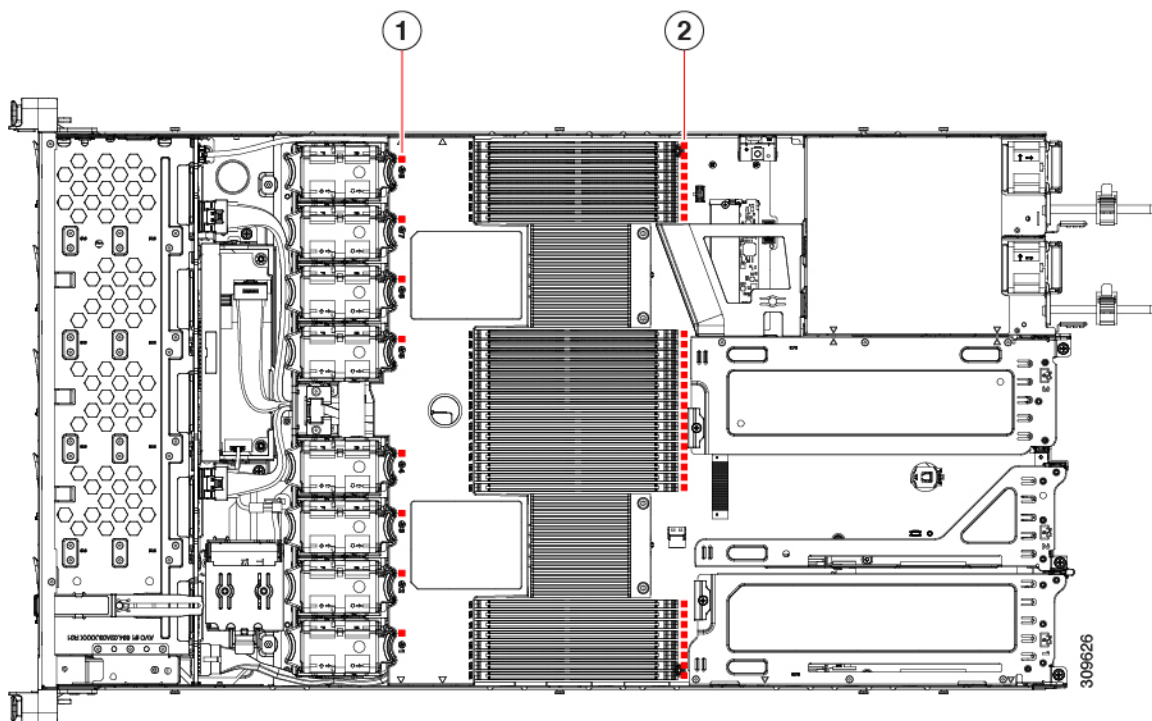
LED Name	States
Rear unit identification	• Off—The unit identification function is not in use. • Blue, blinking—The unit identification function is activated.
2-Gb Ethernet dedicated management link speed	• Off—Link speed is 10 Mbps. • Amber—Link speed is 100 Mbps. • Green—Link speed is 1 Gbps.
3-Gb Ethernet dedicated management link status	• Off—No link is present. • Green—Link is active. • Green, blinking—Traffic is present on the active link.

Power supply status (one LED each power supply unit)	AC power supplies: • Off—No AC input (12 V main power off, 12 V standby power off). • Green, blinking—12 V main power off; 12 V standby power on. • Green, solid—12 V main power on; 12 V standby power on. • Amber, blinking—Warning threshold detected but 12 V main power on. • Amber, solid—Critical error detected; 12 V main power off (for example, over-current, over-voltage, or over-temperature failure). DC power supplies: • Off—No DC input (12 V main power off, 12 V standby power off). • Green, blinking—12 V main power off; 12 V standby power on. • Green, solid—12 V main power on; 12 V standby power on. • Amber, blinking—Warning threshold detected but 12 V main power on. • Amber, solid—Critical error detected; 12 V main power off (for example, over-current, over-voltage, or over-temperature failure).
---	---

Internal Diagnostic LEDs

The server has internal fault LEDs for CPUs, DIMMs, and fan modules.

Figure 22: Internal Diagnostic LED Locations



Fan module fault LEDs (one behind each fan connector on the motherboard)

- Amber—Fan has a fault or is not fully seated.
- Green—Fan is OK.

DIMM fault LEDs (one behind each DIMM socket on the motherboard)

These LEDs operate only when the server is in standby power mode.

- Amber—DIMM has a fault.
- Off—DIMM is OK.

CPU fault LEDs (one behind each CPU socket on the motherboard).

These LEDs operate only when the server is in standby power mode.

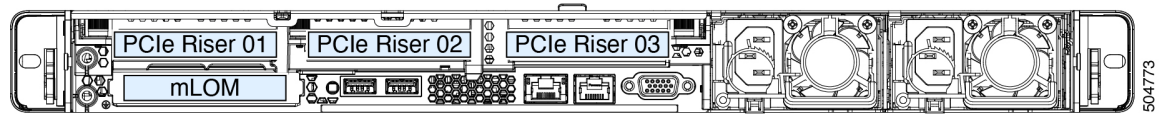
- Amber—CPU has a fault.
- Off—CPU is OK.

-

Physical node cabling

Physical nodes can be deployed in the ND-NODE-L4 (UCS-C225-M6) physical server.

Figure 23: mLOM and PCIe riser 01 card used for node connectivity: ND-NODE-L4 (UCS-C225-M6)



The physical nodes can be deployed with these guidelines:

- All servers come with a Modular LAN on Motherboard (mLOM) card, which you use to connect to the Nexus Dashboard management network.
- The ND-NODE-L4 server includes either a 2x10GbE NIC (APIC-P-ID10Gc), or 2x25/10GbE SFP28 NIC (APIC-P-I8D25GF), or the VIC1455 card in the "PCIe-Riser-01" slot (shown in the above diagram), which you use for Nexus Dashboard data network connectivity.



Note

Verify the port mapping for the data interfaces using the information provided in "Data network connections" before cabling the node. The data interfaces must connect to individual host-facing switch ports; do not configure switch port channels or vPC for these connections.

- Verify that the port-channel feature is enabled on the VIC.
 - For the ND-NODE-L4 (UCS-C225-M6) and ND-NODE-G5S/ND-NODE-G5L (UCS-C225-M8), the port-channel feature is enabled by default at the factory. No additional configurations are necessary for those platforms, but you can use these procedures to verify that the port-channel feature is enabled correctly and enable it on those platforms, if necessary.
 - For the SE-NODE-G2 (UCS-C220-M5), you must manually enable the port-channel feature on the VIC.
1. Log into the CIMC.
 2. Click the hamburger icon, then navigate to:
Networking > Adapter Card 1 > General
 3. In the **Adapter Card Properties** area, locate the **Port Channel** field.
If the **Port Channel** option is disabled (has no check in the check box), enable the feature (put a check in the check box) before beginning the Nexus Dashboard deployment process.

When connecting the node to your management and data networks:

- The interfaces are configured as Linux bonds, one for the data interfaces (bond0) and one for the management interfaces (bond1), running in Active-Standby mode.
- **Management network connections:**
 - You must use the `mgmt0` and `mgmt1` on the mLOM card.
 - All ports must have the same speed, either 1G or 10G.
- **Data network connections:**

- On the ND-NODE-L4 server, you can use the 2x10GbE NIC (APIC-P-ID10GC), or 2x25/10GbE SFP28 NIC (APIC-P-I8D25GF), or the VIC1455 card.



Note If you connect using the 25G Intel NIC, you must disable the FEC setting on the switch port to match the setting on the NIC:

```
(config-if)# fec off
# show interface ethernet 1/34
Ethernet1/34 is up
admin state is up, Dedicated Interface
[...]
FEC mode is off
```

- All ports must have the same speed, either 10G, 25G, or 50G.
- Attach the cables for the data network connections, where one cable connects to `fabric0` and the second cable connects to `fabric1`.

`fabric0` and `fabric1` in the ND-NODE-L4 server corresponds to certain ports, depending on the type of card:

- For the 2-port Intel cards:
 - Port-1 corresponds to `fabric0`
 - Port-2 corresponds to `fabric1`
- For the 4-port cards:
 - Port-1 and Port-2 correspond to `fabric0`
 - Port-3 and Port-4 correspond to `fabric1`

You can therefore have any one of these port channel combinations:

First data cable (<code>fabric0</code>)	Second data cable (<code>fabric1</code>)
Port-1	Port-3
Port-1	Port-4
Port-2	Port-3
Port-2	Port-4

When using a 4-port card on the ND-NODE-L4 server, the order from left to right is Port-4, Port-3, Port-2, Port-1.

You can use both `fabric0` and `fabric1` for data network connectivity as Active-Standby.

**Caution**

- Do not configure a switch port channel or vPC for these connections. The supported pairings listed above identify the valid physical port mappings on the VIC card only.
 - If you connect the two cables for the data network connections using different port channel combinations from those listed above, there will be MAC move notifications on the upstream switch and the ports will flap.
-
- If you connect the nodes to Cisco Catalyst switches, packets are tagged on those Catalyst switches with `vlan0` if no VLAN is specified. In this case, you must add `switchport voice vlan dot1p` command to the switch interfaces where the nodes are connected to ensure reachability over the data network.

What's next

Go to [Deploy Nexus Dashboard as a physical appliance, on page 110](#) to deploy the Nexus Dashboard.

SE-NODE-G2 (UCS-C220-M5)

These sections provide overview and cabling information for the SE-NODE-G2 (UCS-C220-M5) physical appliance.

Understanding the SE-NODE-G2 (UCS-C220-M5)

This section provides overview information for the SE-NODE-G2 (UCS-C220-M5) physical appliance.

The SE-NODE-G2 physical appliance is a Nexus Dashboard appliance that is based on the UCS C220 M5 server. Even though the Nexus Dashboard SE-NODE-G2 is based on the UCS C220 M5 server, you cannot replace components within the Nexus Dashboard SE-NODE-G2 as you would with the base UCS C220 M5 server; instead, you will perform a Return Material Authorization (RMA) operation for the appliance if a component within the Nexus Dashboard SE-NODE-G2 is faulty.

Because the Nexus Dashboard SE-NODE-G2 is based on the UCS C220 M5 server, you can refer to the [Cisco UCS C220 M5 Server Installation and Service Guide](#) for important UCS C220 M5 server information that is also applicable for the Nexus Dashboard SE-NODE-G2, such as the information provided in these chapters in that document:

- Overview
- Installing the Server
- Server Specifications
- Storage Controller Considerations

However, because you cannot replace components within the Nexus Dashboard SE-NODE-G2 as you would with the base UCS C220 M5 server, these chapters from that document are not applicable for the Nexus Dashboard SE-NODE-G2 and should be ignored:

- Maintaining the Server

- GPU Installation
- Installation For Cisco UCS Manager Integration

The following sections provide information that is applicable for the Nexus Dashboard SE-NODE-G2.

Overview

Cisco Nexus Dashboard provides a common platform for deploying Cisco Data Center applications. These applications provide real time analytics, visibility and assurance for policy and infrastructure.

The Cisco Nexus Dashboard server is required for installing and hosting the Cisco Nexus Dashboard application.

The server is orderable in the following version:

- SE-CL-L3 — Small form-factor (SFF) drives, with 10-drive backplane. Supports up to 10 2.5-inch SAS/SATA drives. Drive bays 1 and 2 support NVMe SSDs.

External Features

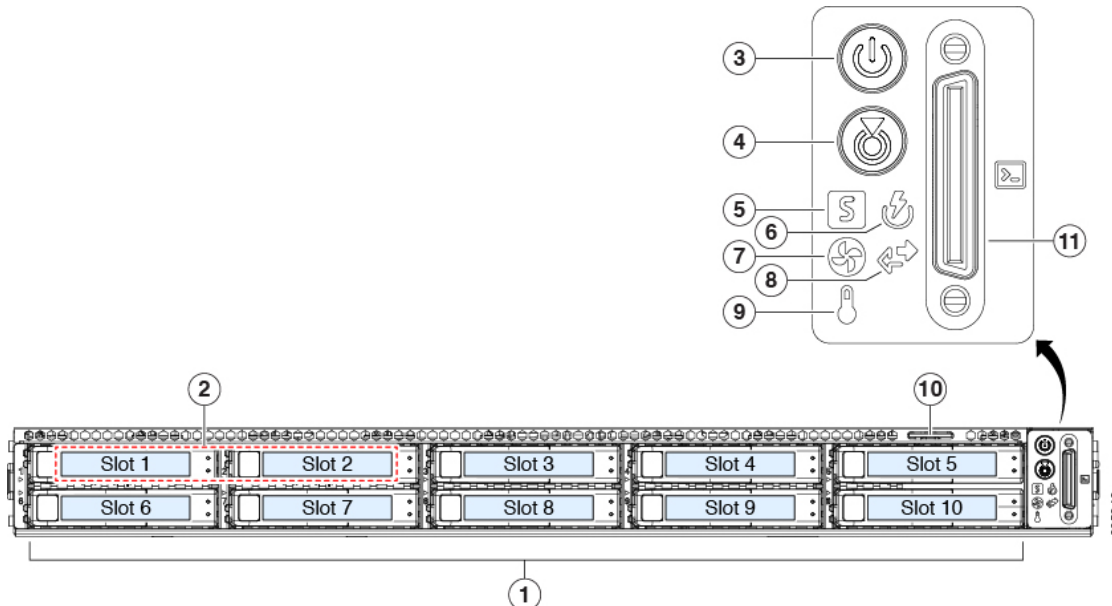
This topic shows the external features of the server versions.

Cisco SE-CL-L3 (SFF Drives) Front Panel Features

The following figure shows the front panel features of the small form-factor drive versions of the server.

For definitions of LED states, see [Front-Panel LEDs, on page 103](#).

Figure 24: Cisco SE-CL-L3 (SFF Drives) Front Panel



Drive bays 1 – 10 support SAS/SATA hard disk drives (HDDs) and solid state drives (SSDs)	Fan status LED
2 • SE-CL-L3 : Drive bays 1 and 2 support NVMe PCIe SSDs.	Network link activity LED

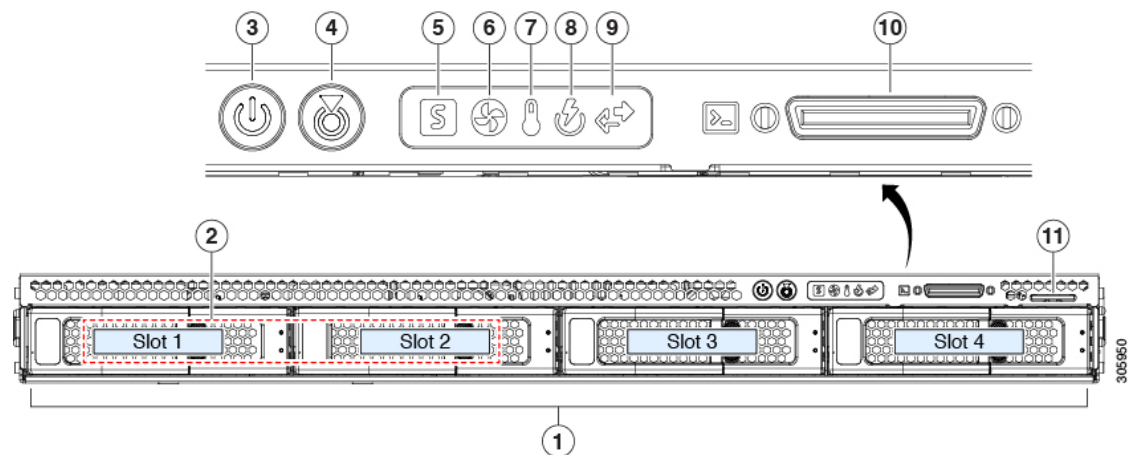
Power button/power status LED	Temperature status LED
Unit identification button/LED	Pull-out asset tag
System status LED	KVM connector (used with KVM cable that provides one DB-15 VGA, one DB-9 serial, and two USB connectors)
Power supply status LED	-

SE-CL-L3 (LFF Drives) Front Panel Features

The following figure shows the front panel features of the large form-factor drive version of the server.

For definitions of LED states, see [Front-Panel LEDs, on page 103](#).

Figure 25: SE-CL-L3 (LFF Drives) Front Panel



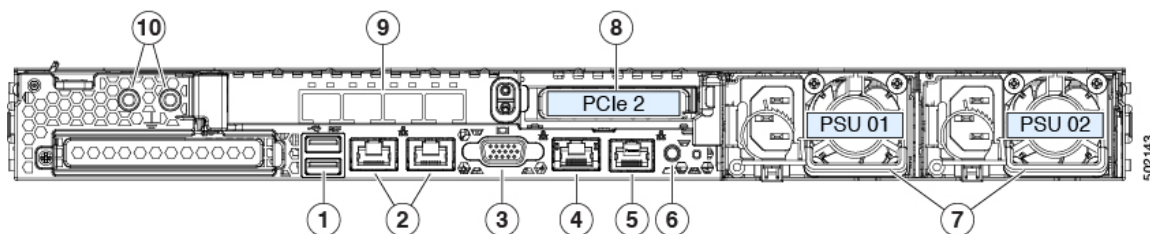
Drive bays 1 – 4 support SAS/SATA HDDs and SSDs	Temperature status LED
Drive bays 1 and 2 support NVMe PCIe SSDs. A size-converter drive sled is required if 2.5-inch SSDs are used.	Power supply status LED
Power button/power status LED	Network link activity LED
Unit identification button/LED	KVM connector (used with KVM cable that provides one DB-15 VGA, one DB-9 serial, and two USB connectors)
System health LED	Pull-out asset tag
Fan status LED	-

SE-CL-L3 Rear Panel Features

The rear panel features are the same for all versions of the server.

For definitions of LED states, see [Rear-Panel LEDs, on page 106](#).

Figure 26: SE-CL-L3 Rear Panel

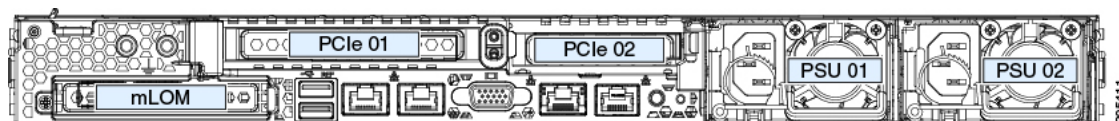


1 USB 3.0 ports (two)	6 Rear unit identification button/LED
2 Dual 1-Gb/10-Gb Ethernet ports (LAN1 and LAN2) The dual LAN ports can support 1 Gbps and 10 Gbps, depending on the link partner capability. These correspond to <code>eth1-1</code> (<code>eth0</code>) and <code>eth1-2</code> (<code>eth1</code>) respectively.	7 Power supplies (two, redundant as 1+1)
3 GA video port (DB-15 connector)	8 PCIe riser 1/slot 1 (x16 lane) Includes PCIe cable connectors for front-loading NVMe SSDs (x8 lane)
4 4-Gb Ethernet dedicated management port	9 Quad 10-Gb/25-Gb ports. These correspond to <code>eth 2-1</code> to <code>eth 2-4</code> . Only 2 interfaces out of the 4 are active at a time (<code>eth2-1/2-2</code> or <code>eth2-3/2-4</code>) in active/standby mode.
5 Serial port (RJ-45 connector)	10 Threaded holes for dual-hole grounding lug

PCIe Slot Specifications

The server contains two PCIe slots on one riser assembly for horizontal installation of PCIe cards. Both slots support the NCSI protocol and 12V standby power.

Figure 27: Rear Panel, Showing PCIe Slot Numbering



The following tables describe the specifications for the slots.

Table 21: PCIe Riser 1/Slot 1

Slot Number	Electrical Lane Width	Connector Length	Maximum Card Length	Card Height (Rear Panel Opening)	NCSI Support
1	Gen-3 x16	x24 connector	$\frac{3}{4}$ length	Full-height	Yes

Micro SD card slot	One socket for Micro SD card
--------------------	------------------------------

Table 22: PCIe Riser 2/Slot 2

Slot Number	Electrical Lane Width	Connector Length	Maximum Card Length	Card Height (Rear Panel Opening)	NCSI Support
2	Gen-3 x16	x24 connector	½ length	½ height	Yes
PCIe cable connector for front-panel NVMe SSDs	Gen-3 x8	Other end of cable connects to front drive backplane to support front-panel NVMe SSDs.			



Note Riser 2/Slot 2 is not available in single-CPU configurations.

Status LEDs and Buttons

This section contains information for interpreting front, rear, and internal LED states.

Front-Panel LEDs

Figure 28: Front Panel LEDs

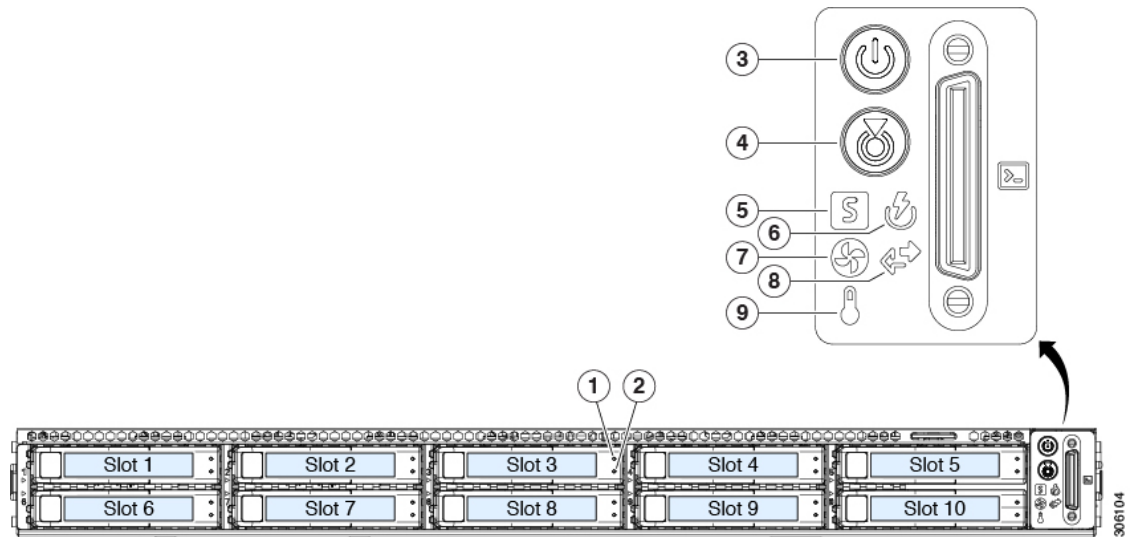


Table 23: Front Panel LEDs, Definition of States

LED Name	States
----------	--------

SAS/SATA drive fault Note NVMe solid state drive (SSD) drive tray LEDs have different behavior than SAS/SATA drive trays.	• Off—The hard drive is operating properly. • Amber—Drive fault detected. • Amber, blinking—The device is rebuilding. • Amber, blinking with one-second interval—Drive locate function activated in the software.
SAS/SATA drive activity LED	• Off—There is no hard drive in the hard drive tray (no access, no fault). • Green—The hard drive is ready. • Green, blinking—The hard drive is reading or writing data.
NVMe SSD drive fault Note NVMe solid state drive (SSD) drive tray LEDs have different behavior than SAS/SATA drive trays.	• Off—The drive is not in use and can be safely removed. • Green—The drive is in use and functioning properly. • Green, blinking—the driver is initializing following insertion or the driver is unloading following an eject command. • Amber—The drive has failed. • Amber, blinking—A drive Locate command has been issued in the software.
NVMe SSD activity	• Off—No drive activity. • Green, blinking—There is drive activity.
Power button/LED	• Off—There is no AC power to the server. • Amber—The server is in standby power mode. Power is supplied only to the Cisco IMC and some motherboard functions. • Green—The server is in main power mode. Power is supplied to all server components.
Unit identification	• Off—The unit identification function is not in use. • Blue, blinking—The unit identification function is activated.

S ystem health	• Green—The server is running in normal operating condition. • Green, blinking—The server is performing system initialization and memory check. • Amber, steady—The server is in a degraded operational state (minor fault). For example: • Power supply redundancy is lost. • CPUs are mismatched. • At least one CPU is faulty. • At least one DIMM is faulty. • At least one drive in a RAID configuration failed. • Amber, 2 blinks—There is a major fault with the system board. • Amber, 3 blinks—There is a major fault with the memory DIMMs. • Amber, 4 blinks—There is a major fault with the CPUs.
P ower supply status	• Green—All power supplies are operating normally. • Amber, steady—One or more power supplies are in a degraded operational state. • Amber, blinking—One or more power supplies are in a critical fault state.
F an status	• Green—All fan modules are operating properly. • Amber, blinking—One or more fan modules breached the non-recoverable threshold.
N etwork link activity	• Off—The Ethernet LOM port link is idle. • Green—One or more Ethernet LOM ports are link-active, but there is no activity. • Green, blinking—One or more Ethernet LOM ports are link-active, with activity.
T emperature status	• Green—The server is operating at normal temperature. • Amber, steady—One or more temperature sensors breached the critical threshold. • Amber, blinking—One or more temperature sensors breached the non-recoverable threshold.

Rear-Panel LEDs

Figure 29: Rear Panel LEDs

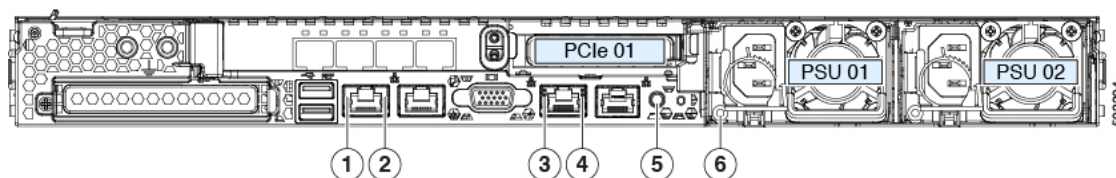


Table 24: Rear Panel LEDs, Definition of States

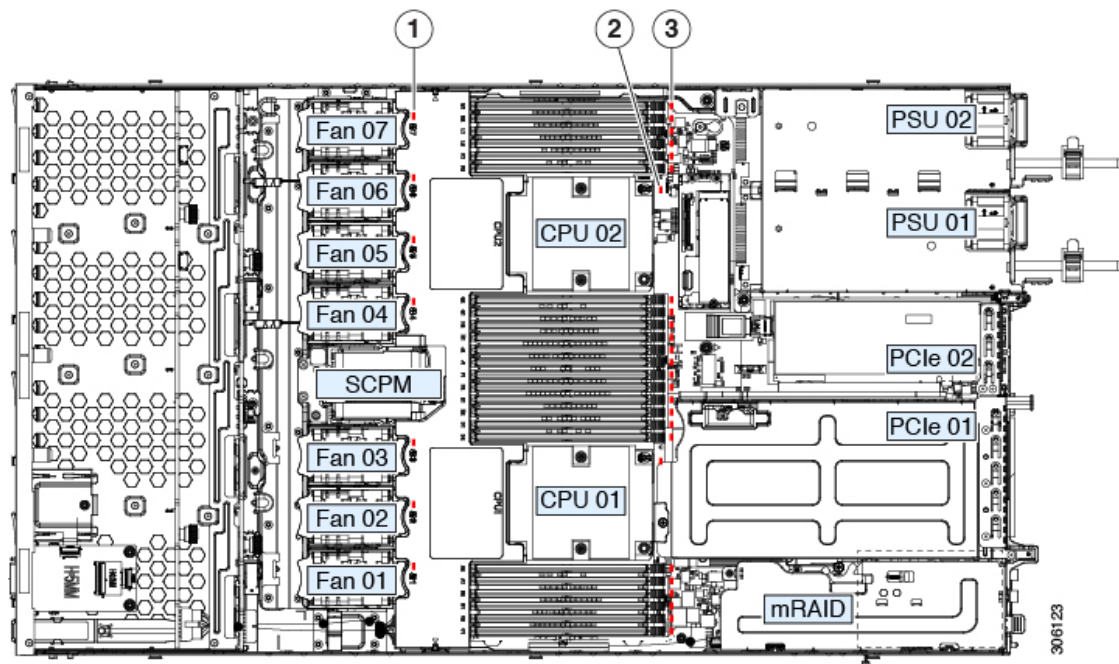
LED Name	States
1-Gb/10-Gb Ethernet link speed (on both LAN1 and LAN2)	• Amber—Link speed is 100 Mbps. • Amber—Link speed is 1 Gbps. • Green—Link speed is 10 Gbps.
2-Gb/10-Gb Ethernet link status (on both LAN1 and LAN2)	• Off—No link is present. • Green—Link is active. • Green, blinking—Traffic is present on the active link.
3-Gb Ethernet dedicated management link speed	• Off—Link speed is 10 Mbps. • Amber—Link speed is 100 Mbps. • Green—Link speed is 1 Gbps.
4-Gb Ethernet dedicated management link status	• Off—No link is present. • Green—Link is active. • Green, blinking—Traffic is present on the active link.
Rear unit identification	• Off—The unit identification function is not in use. • Blue, blinking—The unit identification function is activated.

Power supply status (one LED each power supply unit)	AC power supplies: • Off—No AC input (12 V main power off, 12 V standby power off). • Green, blinking—12 V main power off; 12 V standby power on. • Green, solid—12 V main power on; 12 V standby power on. • Amber, blinking—Warning threshold detected but 12 V main power on. • Amber, solid—Critical error detected; 12 V main power off (for example, over-current, over-voltage, or over-temperature failure). DC power supplies: • Off—No DC input (12 V main power off, 12 V standby power off). • Green, blinking—12 V main power off; 12 V standby power on. • Green, solid—12 V main power on; 12 V standby power on. • Amber, blinking—Warning threshold detected but 12 V main power on. • Amber, solid—Critical error detected; 12 V main power off (for example, over-current, over-voltage, or over-temperature failure).
---	---

Internal Diagnostic LEDs

The server has internal fault LEDs for CPUs, DIMMs, and fan modules.

Figure 30: Internal Diagnostic LED Locations



Fan module fault LEDs (one behind each fan connector on the motherboard)

- Amber—Fan has a fault or is not fully seated.
- Green—Fan is OK.

DIMM fault LEDs (one behind each DIMM socket on the motherboard)

These LEDs operate only when the server is in standby power mode.

- Amber—DIMM has a fault.
- Off—DIMM is OK.

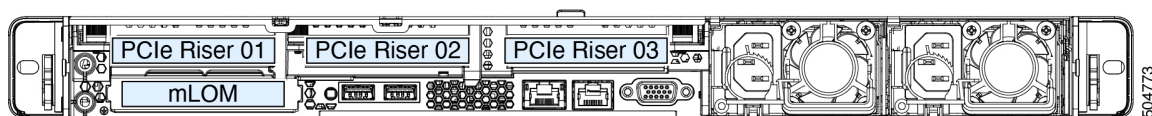
CPU fault LEDs (one behind each CPU socket on the motherboard).

These LEDs operate only when the server is in standby power mode.

- Amber—CPU has a fault.
- Off—CPU is OK.

Physical node cabling

Physical nodes can be deployed in the `SE-NODE-G2` (UCS-C220-M5) physical server.

Figure 31: mLOM and PCIe riser 01 card used for node connectivity: `SE-NODE-G2` (UCS-C220-M5)

The physical nodes can be deployed with these guidelines:

- All servers come with a Modular LAN on Motherboard (mLOM) card, which you use to connect to the Nexus Dashboard management network.
- The SE-NODE-G2 server includes a 4-port VIC1455 card in the "PCIe-Riser-01" slot (shown in the above diagram), which you use for Nexus Dashboard data network connectivity



Note Verify the Cisco UCS VIC card port mapping for the data interfaces using the information provided in "Data network connections" before cabling the node. The data interfaces must connect to individual host-facing switch ports; do not configure switch port channels or vPC for these connections.

- Verify that the port-channel feature is enabled on the VIC.
 - For the ND-NODE-L4 (UCS-C225-M6) and ND-NODE-G5S/ND-NODE-G5L (UCS-C225-M8), the port-channel feature is enabled by default at the factory. No additional configurations are necessary for those platforms, but you can use these procedures to verify that the port-channel feature is enabled correctly and enable it on those platforms, if necessary.
 - For the SE-NODE-G2 (UCS-C220-M5), you must manually enable the port-channel feature on the VIC.
1. Log into the CIMC.
 2. Click the hamburger icon, then navigate to:
Networking > Adapter Card 1 > General
 3. In the **Adapter Card Properties** area, locate the **Port Channel** field.
If the **Port Channel** option is disabled (has no check in the check box), enable the feature (put a check in the check box) before beginning the Nexus Dashboard deployment process.

When connecting the node to your management and data networks:

- The interfaces are configured as Linux bonds, one for the data interfaces (bond0) and one for the management interfaces (bond1), running in Active-Standby mode.
- **Management network connections:**
 - You must use the `mgmt0` and `mgmt1` on the mLOM card.
 - All ports must have the same speed, either 1G or 10G.
- **Data network connections:**
 - On the SE-NODE-G2 server, you must use the VIC1455 card.
 - All ports must have the same speed, either 10G, 25G, or 50G.
 - Attach the cables for the data network connections, where one cable connects to `fabric0` and the second cable connects to `fabric1`.
`fabric0` and `fabric1` in SE-NODE-G2 corresponds to these ports:

- Port-1 and Port-2 correspond to `fabric0`
- Port-3 and Port-4 correspond to `fabric1`

You can therefore have any one of these port channel combinations:

First data cable (<code>fabric0</code>)	Second data cable (<code>fabric1</code>)
Port-1	Port-3
Port-1	Port-4
Port-2	Port-3
Port-2	Port-4

When using a 4-port card on the `SE-NODE-G2` server, the order from left to right is Port-1, Port-2, Port-3, Port-4.

You can use both `fabric0` and `fabric1` for data network connectivity as Active-Standby.



Caution

- Do not configure a switch port channel or vPC for these connections. The supported pairings listed above identify the valid physical port mappings on the VIC card only.
- If you connect the two cables for the data network connections using different port channel combinations from those listed above, there will be MAC move notifications on the upstream switch and the ports will flap.
- If you connect the nodes to Cisco Catalyst switches, packets are tagged on those Catalyst switches with `vlan0` if no VLAN is specified. In this case, you must add `switchport voice vlan dot1p` command to the switch interfaces where the nodes are connected to ensure reachability over the data network.

What's next

Go to [Deploy Nexus Dashboard as a physical appliance, on page 110](#) to deploy the Nexus Dashboard.

Deploy Nexus Dashboard as a physical appliance

When you first receive the Nexus Dashboard physical hardware, it comes preloaded with the software image. Follow these steps to deploy Nexus Dashboard as a physical appliance.

Before you begin

Complete the requirements and guidelines described in [Prerequisites and guidelines for deploying Nexus Dashboard as a physical appliance, on page 63](#).

Procedure

Step 1

Configure the first node's basic information.

You must configure only a single ("first") node as described in this step. Other nodes will be configured during the GUI-based cluster deployment process described in the following steps and will accept settings from the first `primary` node. The other two `primary` nodes do not require any additional configuration besides ensuring that their CIMC IP addresses are reachable from the first `primary` node and login credentials are set, as well as network connectivity between the nodes is established on the data network.

- a) SSH into the node using CIMC management IP and use the `connect host` command to connect to the node's console.

```
C220-WZP23150D4C# connect host
CISCO Serial Over LAN:
Press Ctrl+x to Exit the session
```

After connecting to the host, press **Enter** to continue.

- b) After you see the Nexus Dashboard setup utility prompt, press **Enter**.

```
Starting Nexus Dashboard setup utility
Welcome to Nexus Dashboard 4.2.1
Press Enter to manually bootstrap your first master node...
```

- c) Enter and confirm the `admin` password

This password will be used for the `rescue-user` CLI login as well as the initial GUI password.

```
Admin Password:
Reenter Admin Password:
```

- d) Enter the management network information.

```
Management Network:
IP Address/Mask: 192.168.9.172/24
Gateway: 192.168.9.1
```

Note

If you want to configure IPv6-only mode, enter the IPv6 in the above example instead.

- e) Review and confirm the entered information.

You will be asked if you want to change the entered information. If all the fields are correct, enter the capital letter `N` to proceed. If you want to change any of the entered information, enter `y` to re-start the basic configuration script.

```
Please review the config
Management network:
Gateway: 192.168.9.1
IP Address/Mask: 192.168.9.172/24
```

```
Re-enter config? (y/N): N
```

Step 2

Wait for the process to complete.

After you enter and confirm management network information of the first node, the initial setup configures the networking and brings up the UI, which you will use to add two and configure other nodes and complete the cluster deployment.

```
Please wait for system to boot: [#####] 100%
System up, please wait for UI to be online.
```

System UI online, please login to <https://192.168.9.172> to continue.

Step 3 Open your browser and navigate to <https://<node-mgmt-ip>> to open the GUI.

The rest of the configuration workflow takes place from one of the node's GUI. You can choose any one of the nodes you deployed to begin the bootstrap process and you do not need to log in to or configure the other two nodes directly.

Enter the password you entered in a previous step and click **Login**

Step 4 After logging in, review the **Meet Nexus Dashboard** welcome screens, then proceed to the **Journey** screen. Click on **Go** under the **Cluster bringup** step to begin the bootstrap process.

Step 5 Enter the requested information in the **Basic Information** page of the **Cluster Bringup** wizard.

a) For **Cluster Name**, enter a name for this Nexus Dashboard cluster.

The cluster name must follow the [RFC-1123](#) requirements.

b) For **Select the Nexus Dashboard Implementation type**, choose either **LAN** or **SAN** then click **Next**.

Step 6 Enter the requested information in the **Configuration** page of the **Cluster Bringup** wizard.

- a) (Optional) If you want to enable IPv6 functionality for the cluster, put a check in the **Enable IPv6** checkbox.
- b) Click **+Add DNS provider** to add one or more DNS servers, enter the DNS provider IP address, then click the checkmark icon.
- c) (Optional) Click **+Add DNS search domain** to add a search domain, enter the DNS search domain IP address, then click the checkmark icon.
- d) (Optional) If you want to enable NTP server authentication, put a check in the **NTP Authentication** checkbox.
- e) If you enabled NTP authentication, click **+ Add Key**, enter the required information, and click the checkmark icon to save the information.
 - **Key**—Enter the NTP authentication key, which is a cryptographic key that is used to authenticate the NTP traffic between the Nexus Dashboard and the NTP servers. You will define the NTP servers in the following step, and multiple NTP servers can use the same NTP authentication key.
 - **ID**—Enter a key ID for the NTP host. Each NTP key must be assigned a unique key ID, which is used to identify the appropriate key to use when verifying the NTP packet.
 - **Authentication Type**—Choose authentication type for the NTP key.
 - Put a check in the **Trusted** checkbox if you want this key to be trusted. Untrusted keys cannot be used for NTP authentication.

For the complete list of NTP authentication requirements and guidelines, see [General prerequisites and guidelines, on page 11](#).

If you want to enter additional NTP keys, click **+ Add Key** again and enter the information.

- f) If you enabled NTP authentication, click **+Add NTP Host Name/IP Address**, enter the required information, and click the checkmark icon to save the information.
 - **NTP Host**—Enter an IP address; fully qualified domain names (FQDN) are not supported.
 - **Key ID**—Enter the key ID of the NTP key you defined in the previous substep.
If NTP authentication is disabled, this field is grayed out.
 - Put a check in the **Preferred** checkbox if you want this host to be preferred.

Note

If the node into which you are logged in is configured with only an IPv4 address, but you have checked **Enable IPv6** in a previous step and entered an IPv6 address for an NTP server, you will get the following validation error:

NTP Host*	Key ID	Preferred
2001:420:28e:202a:5054:ff:fe6f:b3f6		true

[+ Add NTP Host Name/IP Address](#)

⚠ Could not validate one or more hosts. Can not reach NTP on Management Network

This is because the node does not have an IPv6 address yet and is unable to connect to an IPv6 address of the NTP server. You will enter IPv6 address in the next step. In this case, enter the other required information as described in the following steps and click **Next** to proceed to the next page where you will enter IPv6 addresses for the nodes.

If you want to enter additional NTP servers, click **+Add NTP Host Name/IP Address** again and enter the information.

- g) For **Proxy Server**, enter the URL or IP address of a proxy server.

For clusters that do not have direct connectivity to Cisco cloud, we recommend configuring a proxy server to establish the connectivity. This allows you to mitigate risk from exposure to non-conformant hardware and software in your fabrics.

You can click **+Add Ignore Host** to enter one or more destination IP addresses for which traffic will skip using the proxy.

The proxy server must permit these URLs:

```
svc.intersight.com
svc-static1.intersight.com
svc-static1.ucs-connect.com
```

If you do not want to configure a proxy, click **Skip Proxy** then click **Confirm**.

- h) (Optional) If your proxy server requires authentication, put a check in the **Authentication required for Proxy** checkbox and enter the login credentials.
- i) (Optional) Expand the **Advanced Settings** category and change the settings if required.

Under advanced settings, you can configure these settings:

- **App Network**—The address space used by the application's services running in the Nexus Dashboard. Enter the IP address and netmask.
- **Service Network**—An internal network used by Nexus Dashboard and its processes. Enter the IP address and netmask.
- **App Network IPv6**—If you put a check in the **Enable IPv6** checkbox earlier, enter the IPv6 subnet for the app network.
- **Service Network IPv6**—If you put a check in the **Enable IPv6** checkbox earlier, enter the IPv6 subnet for the service network.

For more information about the application and service networks, see [General prerequisites and guidelines, on page 11](#).

- j) Click **Next**.

Step 7

In the **Node Details** page, update the first node's information.

You have defined the Management network and IP address for the node into which you are currently logged in during the initial node configuration in earlier steps, but you must also enter the Data network information for the node before you can proceed with adding the other `primary` nodes and creating the cluster.

- a) For **Cluster Connectivity**, if your cluster is deployed in L3 mode, choose **BGP**. Otherwise, choose **L2**.

BGP configuration is required for the persistent IP addresses feature used by telemetry. This feature is described in more detail in the [BGP configuration and persistent IP addresses, on page 46](#) and [Nexus Dashboard persistent IP addresses, on page 40](#) sections.

Note

You can enable BGP at this time or in the Nexus Dashboard GUI after the cluster is deployed. All remaining nodes need to configure BGP if it is configured. You must enable BGP now if the data network of nodes have different subnets.

- b) Click the **Edit** button next to the first node.

The node's **Serial Number**, **Management Network** information, and **Type** are automatically populated, but you must enter the other information.

- c) For **Name**, enter a name for the node.

The node's **Name** will be set as its hostname, so it must follow the [RFC-1123](#) requirements.

Note

If you need to change the name but the **Name** field is not editable, run the CIMC validation again to fix this issue.

- d) For **Type**, choose **Primary**.

The first nodes of the cluster must be set to **Primary**. You will add the secondary nodes in a later step if required for higher scale.

- e) In the **Data Network** area, enter the node's data network information.

Enter the data network IP address, netmask, and gateway. Optionally, you can also enter the VLAN ID for the network. Leave the VLAN ID field blank if your configuration does not require VLAN. If you chose **BGP** for **Cluster Connectivity**, enter the ASN.

If you enabled IPv6 functionality in a previous page, you must also enter the IPv6 address, netmask, and gateway.

Note

If you want to enter IPv6 information, you must do so during the cluster bootstrap process. To change the IP address configuration later, you would need to redeploy the cluster.

All nodes in the cluster must be configured with either only IPv4, only IPv6, or dual stack IPv4/IPv6.

- f) If you chose **BGP** for **Cluster Connectivity**, then in the **BGP peer details** area, enter the peer's IPv4 address and ASN.

You can click + **Add IPv4 BGP peer** to add additional peers.

If you enabled IPv6 functionality in a previous page, you must also enter the peer's IPv6 address and ASN.

- g) Click **Save** to save the changes.

Step 8

If you are deploying a multi-node cluster, in the **Node Details** page, click **Add Node** to add the second node to the cluster.

- a) In the **Deployment Details** area, enter the **CIMC IP Address**, **Username**, and **Password** for the second node.

Note

For **Username** for the second node, enter the admin user ID.

- b) Click **Validate** to verify connectivity to the node.

The node's serial number is automatically populated after CIMC connectivity is validated.

- c) For **Name**, enter the name for the node.

The node's name will be set as its hostname, so it must follow the [RFC-1123](#) requirements.

- d) For **Type**, choose `Primary`.

The first 3 nodes of the cluster must be set to `Primary`. You will add the secondary nodes in a later step if required for higher scale.

- e) In the **Management Network** area, enter the node's management network information.

You must enter the management network IP address, netmask, and gateway.

If you enabled IPv6 functionality in a previous page, you must also enter the IPv6 address, netmask, and gateway.

Note

All nodes in the cluster must be configured with either only IPv4, only IPv6, or dual stack IPv4/IPv6.

- f) In the **Data Network** area, enter the node's data network information.

Enter the data network IP address, netmask, and gateway. Optionally, you can also enter the VLAN ID for the network. Leave the VLAN ID field blank if your configuration does not require VLAN. If you chose **BGP for Cluster Connectivity**, enter the ASN.

If you enabled IPv6 functionality in a previous page, you must also enter the IPv6 address, netmask, and gateway.

Note

If you want to enter IPv6 information, you must do so during the cluster bootstrap process. To change the IP address configuration later, you would need to redeploy the cluster.

All nodes in the cluster must be configured with either only IPv4, only IPv6, or dual stack IPv4 and IPv6.

- g) If you chose **BGP for Cluster Connectivity**, then in the **BGP peer details** area, enter the peer's IPv4 address and ASN.

You can click + **Add IPv4 BGP peer** to add additional peers.

If you enabled IPv6 functionality in a previous page, you must also enter the peer's IPv6 address and ASN.

- h) Click **Save** to save the changes.

- i) Repeat this step for the final (third) primary node of the cluster.

Step 9

(Optional) Repeat the previous step to enter information about any additional secondary or standby nodes.

Note

To support higher scale, you must provide a sufficient number of secondary nodes during deployment. Refer to the [Nexus Dashboard Cluster Sizing](#) tool for exact number of additional secondary nodes required for your specific use case.

You can choose to add the standby nodes now or at a later time after the cluster is deployed.

Step 10

In the **Node Details** page, verify the information that you entered, then click **Next**.

Step 11 In the **Persistent IPs** page, if you want to add more persistent IP addresses, click + **Add Data Service IP Address**, enter the IP address, and click the checkmark icon. Repeat this step as many times as desired, then click **Next**.

You must configure the minimum number of required persistent IP addresses during the bootstrap process. This step enables you to add more persistent IP addresses if desired.

Step 12 In the **Summary** page, review and verify the configuration information, click **Save**, and click **Continue** to confirm the correct deployment mode and proceed with building the cluster.

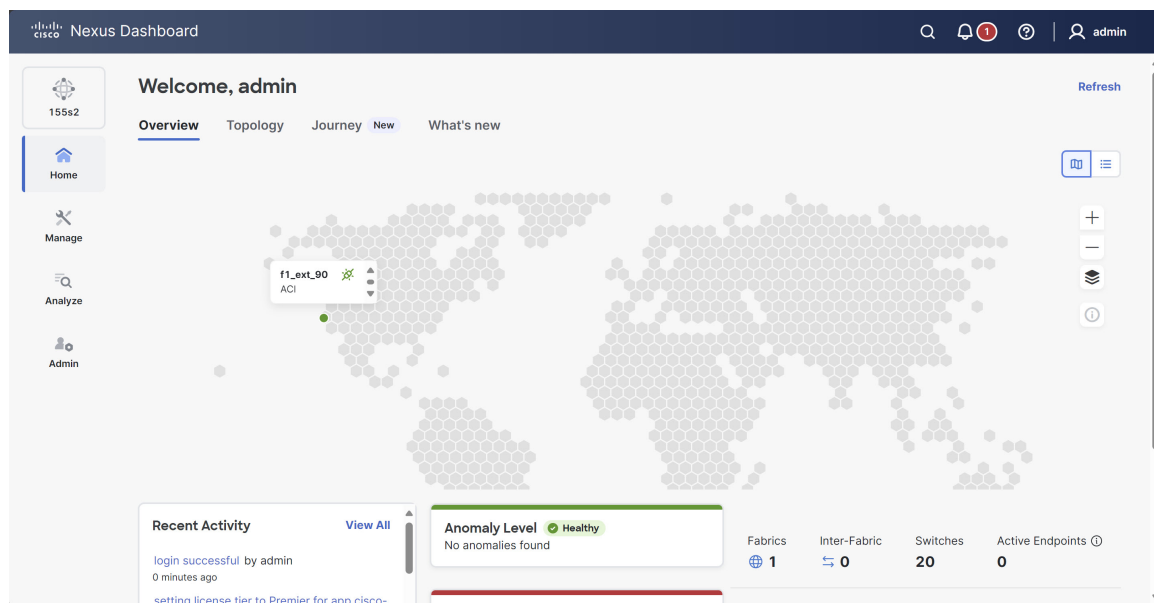
During the node bootstrap and cluster bring-up, the overall progress as well as each node's individual progress will be displayed in the UI. If you do not see the bootstrap progress advance, manually refresh the page in your browser to update the status.

It may take up to 60 minutes or more for the cluster to form, depending on the number of nodes in the cluster, and all the features to start. When cluster configuration is complete, the page will reload to the Nexus Dashboard GUI.

Step 13 Verify that the cluster is healthy.

After the cluster becomes available, you can access it by browsing to any one of your nodes' management IP addresses. The default password for the `admin` user is the same as the `rescue-user` password you chose for the first node. During this time, the UI will display a banner at the top stating "Service Installation is in progress, Nexus Dashboard configuration tasks are currently disabled".

After all the cluster is deployed and all services are started, you can look at the **Anomaly Level** on the **Home > Overview** page to ensure the cluster is healthy:



Alternatively, you can log in to any one node using SSH as the `rescue-user` using the password you entered during node deployment and using the `acs health` command to see the status:

- While the cluster is converging, you may see the following output:

```
$ acs health
k8s install is in-progress

$ acs health
k8s services not in desired state - [...]
```

```
$ acs health  
k8s: Etcd cluster is not ready
```

- When the cluster is up and running, the following output will be displayed:

```
$ acs health  
All components are healthy
```

Note

In some situations, you might power cycle a node (power it off and then back on) and find it stuck in this stage:

```
deploy base system services
```

This is due to an issue with `etcd` on the node after a reboot of the physical Nexus Dashboard cluster.

To resolve the issue, enter the `acs reboot clean` command on the affected node.

Step 14 (Optional) Connect your Cisco Nexus Dashboard cluster to Cisco Intersight for added visibility and benefits. Refer to [Working with Cisco Intersight](#) for detailed steps.

Step 15 After you have deployed Nexus Dashboard, see the [collections page](#) for this release for configuration information.

What to do next

The next task is to create the fabrics and fabric groups. See the *Creating Fabrics and Fabric Groups* article for this release on the [Cisco Nexus Dashboard collections page](#).



CHAPTER 6

Deploying in VMware ESX

- [Prerequisites and guidelines for deploying the Nexus Dashboard cluster in VMware ESX, on page 119](#)
- [Deploy Nexus Dashboard Using VMware vCenter, on page 121](#)
- [Deploy Nexus Dashboard Directly in VMware ESXi, on page 135](#)

Prerequisites and guidelines for deploying the Nexus Dashboard cluster in VMware ESX

Before you proceed with deploying the Nexus Dashboard cluster in VMware ESX, you must:

- Ensure that the ESX form factor supports your scale requirements.

Scale support and co-hosting vary based on the cluster form factor you plan to deploy. You can use the [Nexus Dashboard Capacity Planning](#) tool to verify that the virtual form factor satisfies your deployment requirements.



Note

Some deployments may require only a single ESX virtual node for one or more specific use cases. In that case, the capacity planning tool will indicate the requirement and you can simply skip the additional node deployment step in the following sections.

- Review and complete the general prerequisites described in [Prerequisites and Guidelines, on page 11](#).

This document describes how to initially deploy the base Nexus Dashboard cluster. If you want to expand an existing cluster with additional nodes (such as `secondary` or `standby`), see the "Infrastructure Management" chapter of the *Cisco Nexus Dashboard User Guide* instead, which is available from the Nexus Dashboard UI or online at [Cisco Nexus Dashboard User Guide](#)

- Ensure that the CPU family used for the Nexus Dashboard VMs supports AVX instruction set.
- Choose the type of node to deploy:
 - Data node—Node profile with higher system requirements designed for specific Nexus Dashboard features that require the additional resources.
 - App node—Node profile with a smaller resource footprint that can be used for most Nexus Dashboard features.



Note Some larger scale deployments may require additional secondary nodes. If you plan to add secondary nodes to your Nexus Dashboard cluster, you can deploy all nodes (the initial 3-node cluster and the additional secondary nodes) using the OVA-App profile. Detailed scale information is available in the [Cisco Nexus Dashboard Verified Scalability Guide](#) for your release.

Ensure you have enough system resources:

Table 25: Deployment requirements

Data node requirements	App node requirements
- VMware ESXi 7.0, 7.0.1, 7.0.2, 7.0.3, 8.0, 8.0.2, 8.0.3, 9.0.2 - VMware vCenter 7.0.1, 7.0.2, 7.0.3, 8.0, 8.0.2, 8.0.3, 9.0.2 if deploying using VMware vCenter - Each node/VM requires the following: 32 vCPUs with physical CPU reservation of at least 35,200 MHz 128GB of RAM with physical reservation 3TB SSD storage for the data volume and an additional 50GB for the system volume Data nodes must be deployed on storage with the following minimum performance requirements: - The SSD must be attached to the data store directly or in JBOD mode if using a RAID Host Bus Adapter (HBA) - The SSDs must be optimized for Mixed Use/Application (not Read-Optimized) 4K Random Read IOPS: 93000 4K Random Write IOPS: 31000 - We recommend that each Nexus Dashboard node is deployed in a different ESXi server.	- VMware ESXi 7.0, 7.0.1, 7.0.2, 7.0.3, 8.0, 8.0.2, 8.0.3, 9.0.2 - VMware vCenter 7.0.1, 7.0.2, 7.0.3, 8.0, 8.0.2, 8.0.3, 9.0.2 if deploying using VMware vCenter - Each node/VM requires the following: 16 vCPUs with physical CPU reservation of at least 17,600 MHz 64GB of RAM with physical reservation 500GB HDD or SSD storage for the data volume and an additional 50GB for the system volume Some features require App nodes to be deployed on faster SSD storage while other features support HDD. Check the Nexus Dashboard Capacity Planning tool to ensure that you use the correct type of storage. - We recommend that each Nexus Dashboard node is deployed in a different ESXi server.

- If you plan to configure VLAN ID for the cluster's data interfaces, you must enable VLAN 4095 on the data interface port group in VMware vCenter for Virtual Guest VLAN Tagging (VGT) mode. If you

specify a VLAN ID for Nexus Dashboard data interfaces, the packets must carry a Dot1q tag with that VLAN ID. When you set an explicit VLAN tag in a port group in the vSwitch and attach it to a Nexus Dashboard VM's VNIC, the vSwitch removes the Dot1q tag from the packet coming from the uplink before it sends the packet to that VNIC. Because the virtual Nexus Dashboard node expects the Dot1q tag, you must enable VLAN 4095 on the data interface port group to allow all VLANs.

- After each node's VM is deployed, ensure that the VMware Tools' periodic time synchronization is disabled as described in the deployment procedure in the next section.
 - VMware vMotion is not supported for Nexus Dashboard cluster nodes.
 - VMware Distributed Resource Scheduler (DRS) is not supported for Nexus Dashboard cluster nodes.
- If you have DRS enabled at the ESXi cluster level, you must explicitly disable it for the Nexus Dashboard VMs during deployment as described in the following section.
- Deploying using the content library is not supported.
 - VMware snapshots are supported only for Nexus Dashboard VMs that are powered off and must be done for all Nexus Dashboard VMs belonging to the same cluster.

Snapshots of powered on VMs are not supported.

- Cisco does not support the use of nested virtualization environments. Deploying Nexus Dashboard on a virtual machine that is itself running on a virtualized hypervisor (for example, KVM on ESXi) is not a supported configuration and may result in performance degradation or system instability.
- You can choose to deploy the nodes directly in ESXi or using VMware vCenter.

If you want to deploy using VMware vCenter, following the steps described in [Deploy Nexus Dashboard Using VMware vCenter, on page 121](#).

If you want to deploy directly in ESXi, following the steps described in [Deploy Nexus Dashboard Directly in VMware ESXi, on page 135](#).

Deploy Nexus Dashboard Using VMware vCenter

This section describes how to deploy Cisco Nexus Dashboard cluster using VMware vCenter. If you prefer to deploy directly in ESXi, follow the steps described in [Deploy Nexus Dashboard Directly in VMware ESXi, on page 135](#) instead.

Before you begin

- Ensure that you meet the requirements and guidelines described in [Prerequisites and guidelines for deploying the Nexus Dashboard cluster in VMware ESX, on page 119](#).

Procedure

Step 1

Obtain the Cisco Nexus Dashboard OVA image.

- a) Browse to the Software Download page.

<https://software.cisco.com/download/home/286327743/type/286328258/>

- b) Choose the Nexus Dashboard release version you want to download.
- c) Click the **Download** icon next to the Nexus Dashboard OVA image (nd-dk9.<version>.ova).

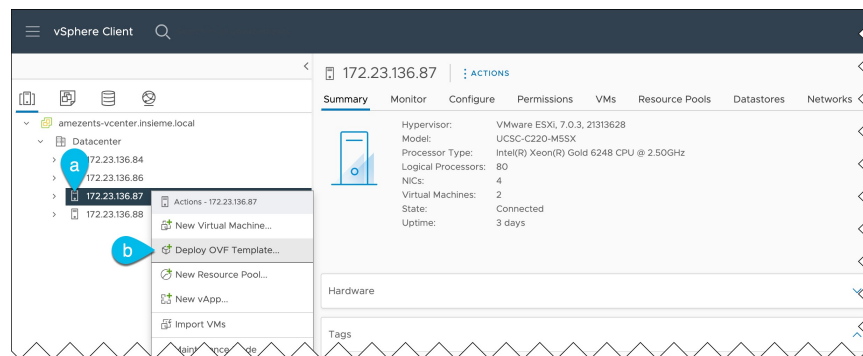
Step 2

Log in to your VMware vCenter.

Depending on the version of your vSphere client, the location and order of configuration screens may differ slightly. The following steps provide deployment details using VMware vSphere Client 7.0.

Step 3

Start the new VM deployment.

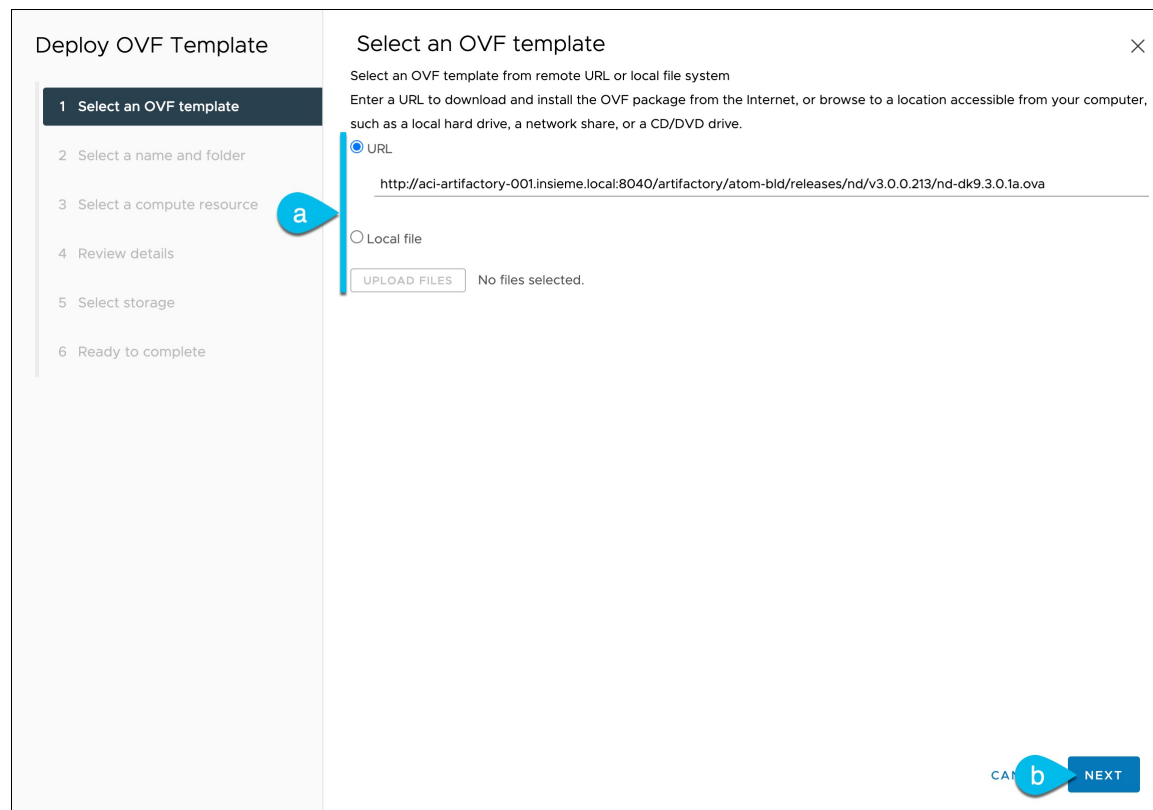


- a) Right-click the ESX host where you want to deploy the VM.
- b) Select **Deploy OVF Template...**

The **Deploy OVF Template** wizard appears.

Step 4

In the **Select an OVF template** screen, provide the OVA image.



- a) Provide the location of the image.

If you hosted the image on a web server in your environment, select **URL** and provide the URL to the image as shown in the above screenshot.

If your image is local, select **Local file** and click **Choose Files** to select the OVA file you downloaded.

- b) Click **Next** to continue.

Step 5

In the **Select a name and folder** screen, provide a name and location for the VM.

The screenshot shows the 'Deploy OVF Template' wizard with six steps. Step 2, 'Select a name and folder', is active. The main area has two sections: 'Specify a unique name and target location' and 'Select a location for the virtual machine.' The first section has a text input for 'Virtual machine name:' with the value 'nd-ova-node1' and a blue callout 'a' pointing to it. The second section has a dropdown menu with 'Datacenter' selected and a blue callout 'b' pointing to it. At the bottom right are 'CANCEL', 'BACK', and 'NEXT' buttons, with a blue callout 'c' pointing to the 'NEXT' button.

- a) Provide the name for the virtual machine.

For example, `nd-ova-node1`.

- b) Select the location for the virtual machine.

- c) Click **Next** to continue

Step 6

In the **Select a compute resource** screen, select the ESX host.

Deploy OVF Template

- 1 Select an OVF template
- 2 Select a name and folder
- 3 Select a compute resource
- 4 Review details
- 5 Select storage
- 6 Ready to complete

Select a compute resource

Select the destination compute resource for this operation

- ✓ Datacenter
 - > 172.23.136.84
 - > 172.23.136.86
 - > 172.23.136.87
 - > 172.23.136.88

Compatibility

✓ Compatibility checks succeeded.

CANCEL BACK NEXT

a) Select the vCenter data center and the ESX host for the virtual machine.

b) Click **Next** to continue

Step 7

In the **Review details** screen, click **Next** to continue.

Step 8

In the **Configuration** screen, select the node profile you want to deploy.

The screenshot shows the 'Deploy OVF Template' wizard in the VMware vSphere Client. The left sidebar contains a list of steps: 1. Select an OVF template, 2. Select a name and folder, 3. Select a compute resource, 4. Review details, 5. Configuration (highlighted), 6. Select storage, 7. Select networks, 8. Customize template, and 9. Ready to complete. The main area is titled 'Configuration' and contains a section 'Select a deployment configuration' with two radio buttons: 'Data' (selected) and 'App'. To the right of these buttons is a 'Description' box that reads: 'Use this deployment profile to configure an OVA with 32 vCPUs, 128 GB RAM, and 3 TB Disk. Read https://www.cisco.com/c/dam/en/us/td/c/sizing/index.html for information on capacity planning'. At the bottom right of the main area, there are three buttons: 'CANCEL', 'BACK', and 'NEXT'.

Configuration	Description
☒ Data	Use this deployment profile to configure an OVA with 32 vCPUs, 128 GB RAM, and 3 TB Disk. Read https://www.cisco.com/c/dam/en/us/td/c/sizing/index.html for information on capacity planning
☐ App	

2 Items

CANCEL BACK NEXT

- a) Select either `App` or `Data` node profile based on your use case requirements.

For more information about the node profiles, see [Prerequisites and guidelines for deploying the Nexus Dashboard cluster in VMware ESX, on page 119](#).

- b) Click **Next** to continue

Step 9

In the **Select storage** screen, provide the storage information.

Deploy OVF Template

- Select an OVF template
- Select a name and folder
- Select a compute resource
- Review details
- Configuration
- Select storage**
- Select networks
- Customize template
- Ready to complete

Select storage

Select the storage for the configuration and disk files

☐ Encrypt this virtual machine (Requires Key Management Server)

Select virtual disk format: Thick Provision Lazy Zeroed

VM Storage Policy: Datastore Default

☒ Disable Storage DRS for this virtual machine

	Name	Storage Compatibility	Capacity	Provisioned	Free	Type	Cluster
☐	datastore1	--	989.75 GB	613.47 GB	376.28 GB	VMFS 6	
☒	datastore2-s...	--	3.49 TB	1.55 TB	1.94 TB	VMFS 6	
☐	datastore3-s...	--	3.49 TB	1.46 GB	3.49 TB	VMFS 6	
☐	datastore4-s...	--	3.49 TB	1.46 GB	3.49 TB	VMFS 6	

4 Items

Compatibility

✓ Compatibility checks succeeded.

CANCEL BACK NEXT

a) Select the datastore for the virtual machine.

We recommend a unique datastore for each node.

b) Check the **Disable Storage DRS for this virtual machine** checkbox.

Nexus Dashboard does not support VMware DRS.

c) From the **Select virtual disk format** drop-down, choose `Thick Provisioning Lazy Zeroed`.

d) Click **Next** to continue

Step 10

In the **Select networks** screen, choose the VM network for the Nexus Dashboard's Management and Data networks and click **Next** to continue.

There are two networks required by the Nexus Dashboard cluster, both of which that have ports configured for high availability:

- **Management network:** The bonded ports **mgmt0/mgmt1** are used for the Nexus Dashboard cluster's management network.
- **Data network:** The bonded ports **fabric0/fabric1** are used for the Nexus Dashboard cluster's data network.

For more information about these networks, see [General prerequisites and guidelines, on page 11](#) in the "Deployment Overview and Requirements" chapter.

Step 11

In the **Customize template** screen, provide the required information.

Deploy OVF Template

- Select an OVF template
- Select a name and folder
- Select a compute resource
- Review details
- Configuration
- Select storage
- Select networks
- Customize template**
- Ready to complete

Customize template

Customize the deployment properties of this software solution.

☒ All properties have valid values

Node Configuration		4 settings
1. Password	Local "rescue-user" password	
Password	*****	
Confirm Password	*****	
2. Management Network IP Address	Management network IP address only. Examples: 192.168.1.100 or 2222::32 172.23.52.121	
3. Management Network Subnet Mask	CIDR prefix length only (no dotted mask). Use 1-32 for IPv4 (e.g. 24) or up to 128 for IPv6 (e.g. 64, 120) 21	
4. Management Gateway IP	Management network gateway IP address. Enter IP only Ex: 192.168.1.1 or 2222::1 172.23.48.1	

CANCEL
BACK
NEXT

- a) Provide and confirm the **Password**.

This password is used for the `rescue-user` account on each node.

Note

You must provide the same password for all nodes or the cluster creation will fail.

- b) Provide the **Management Network IP Address**.

For example, `192.168.1.100`.

- c) Provide the **Management Network Subnet Mask**.

This entry must be a CIDR prefix length only. For example:

- 24 or 32 for IPv4
- Up to 128 for IPv6

- d) Provide the **Management Gateway IP**.

- e) Click **Next** to continue.

Step 12 In the **Ready to complete** screen, verify that all information is accurate and click **Finish** to begin deploying the first node.

Step 13 Repeat previous steps to deploy the additional nodes.

Note

If you are deploying a single-node cluster, you can skip this step.

For multi-node clusters, you must deploy two additional `Primary` nodes and as many `Secondary` nodes as required by your specific use case. The total number of required nodes is available in the [Nexus Dashboard Capacity Planning](#) tool.

You do not need to wait for the first node's VM deployment to complete, you can begin deploying the other two nodes simultaneously. The steps to deploy the second and third nodes are identical to the first node's.

Step 14 Wait for the VM(s) to finish deploying.

Step 15 Ensure that the VMware Tools periodic time synchronization is disabled, then start the VMs.

To disable time synchronization:

- a) Right-click the node's VM and select **Edit Settings**.
- b) In the **Edit Settings** window, select the **VM Options** tab.
- c) Expand the **VMware Tools** category and uncheck the **Synchronize time periodically** option.

Step 16 Open your browser and navigate to `https://<node-mgmt-ip>` to open the GUI.

The rest of the configuration workflow takes place from one of the node's GUI. You can choose any one of the nodes you deployed to begin the bootstrap process and you do not need to log in to or configure the other two nodes directly.

Enter the password you entered in a previous step and click **Login**

Step 17 After logging in, review the **Meet Nexus Dashboard** welcome screens, then proceed to the **Journey** screen. Click on **Go** under the **Cluster bringup** step to begin the bootstrap process.

Step 18 Enter the requested information in the **Basic Information** page of the **Cluster Bringup** wizard.

- a) For **Cluster Name**, enter a name for this Nexus Dashboard cluster.

The cluster name must follow the [RFC-1123](#) requirements.

- b) For **Select the Nexus Dashboard Implementation type**, choose either **LAN** or **SAN** then click **Next**.

Step 19 Enter the requested information in the **Configuration** page of the **Cluster Bringup** wizard.

- a) (Optional) If you want to enable IPv6 functionality for the cluster, put a check in the **Enable IPv6** checkbox.
- b) Click **+Add DNS provider** to add one or more DNS servers, enter the DNS provider IP address, then click the checkmark icon.
- c) (Optional) Click **+Add DNS search domain** to add a search domain, enter the DNS search domain IP address, then click the checkmark icon.
- d) (Optional) If you want to enable NTP server authentication, put a check in the **NTP Authentication** checkbox.
- e) If you enabled NTP authentication, click **+ Add Key**, enter the required information, and click the checkmark icon to save the information.

- **Key**—Enter the NTP authentication key, which is a cryptographic key that is used to authenticate the NTP traffic between the Nexus Dashboard and the NTP servers. You will define the NTP servers in the following step, and multiple NTP servers can use the same NTP authentication key.
- **ID**—Enter a key ID for the NTP host. Each NTP key must be assigned a unique key ID, which is used to identify the appropriate key to use when verifying the NTP packet.
- **Authentication Type**—Choose authentication type for the NTP key.
- Put a check in the **Trusted** checkbox if you want this key to be trusted. Untrusted keys cannot be used for NTP authentication.

For the complete list of NTP authentication requirements and guidelines, see [General prerequisites and guidelines, on page 11](#).

If you want to enter additional NTP keys, click **+ Add Key** again and enter the information.

- f) If you enabled NTP authentication, click **+Add NTP Host Name/IP Address**, enter the required information, and click the checkmark icon to save the information.

- **NTP Host**—Enter an IP address; fully qualified domain names (FQDN) are not supported.
- **Key ID**—Enter the key ID of the NTP key you defined in the previous substep.
If NTP authentication is disabled, this field is grayed out.
- Put a check in the **Preferred** checkbox if you want this host to be preferred.

Note

If the node into which you are logged in is configured with only an IPv4 address, but you have checked **Enable IPv6** in a previous step and entered an IPv6 address for an NTP server, you will get the following validation error:

NTP Host*	Key ID	Preferred	
2001:420:28e:202a:5054:ff:fe6f:b3f6		true	 
 Add NTP Host Name/IP Address			

 Could not validate one or more hosts Can not reach NTP on Management Network

This is because the node does not have an IPv6 address yet and is unable to connect to an IPv6 address of the NTP server. You will enter IPv6 address in the next step. In this case, enter the other required information as described in the following steps and click **Next** to proceed to the next page where you will enter IPv6 addresses for the nodes.

If you want to enter additional NTP servers, click **+Add NTP Host Name/IP Address** again and enter the information.

- g) For **Proxy Server**, enter the URL or IP address of a proxy server.

For clusters that do not have direct connectivity to Cisco cloud, we recommend configuring a proxy server to establish the connectivity. This allows you to mitigate risk from exposure to non-conformant hardware and software in your fabrics.

You can click **+Add Ignore Host** to enter one or more destination IP addresses for which traffic will skip using the proxy.

The proxy server must permit these URLs:

```
svc.intersight.com
svc-static1.intersight.com
svc-static1.ucs-connect.com
```

If you do not want to configure a proxy, click **Skip Proxy** then click **Confirm**.

- h) (Optional) If your proxy server requires authentication, put a check in the **Authentication required for Proxy** checkbox and enter the login credentials.
- i) (Optional) Expand the **Advanced Settings** category and change the settings if required.

Under advanced settings, you can configure these settings:

- **App Network**—The address space used by the application's services running in the Nexus Dashboard. Enter the IP address and netmask.
- **Service Network**—An internal network used by Nexus Dashboard and its processes. Enter the IP address and netmask.
- **App Network IPv6**—If you put a check in the **Enable IPv6** checkbox earlier, enter the IPv6 subnet for the app network.

- **Service Network IPv6**—If you put a check in the **Enable IPv6** checkbox earlier, enter the IPv6 subnet for the service network.

For more information about the application and service networks, see [General prerequisites and guidelines, on page 11](#).

j) Click **Next**.

Step 20

In the **Node Details** page, update the first node's information.

You have defined the Management network and IP address for the node into which you are currently logged in during the initial node configuration in earlier steps, but you must also enter the Data network information for the node before you can proceed with adding the other `primary` nodes and creating the cluster.

- a) For **Cluster Connectivity**, if your cluster is deployed in L3 mode, choose **BGP**. Otherwise, choose **L2**.

BGP configuration is required for the persistent IP addresses feature used by telemetry. This feature is described in more detail in the [BGP configuration and persistent IP addresses, on page 46](#) and [Nexus Dashboard persistent IP addresses, on page 40](#) sections.

Note

You can enable BGP at this time or in the Nexus Dashboard GUI after the cluster is deployed. All remaining nodes need to configure BGP if it is configured. You must enable BGP now if the data network of nodes have different subnets.

- b) Click the **Edit** button next to the first node.

The node's **Serial Number**, **Management Network** information, and **Type** are automatically populated, but you must enter the other information.

- c) For **Name**, enter a name for the node.

The node's **Name** will be set as its hostname, so it must follow the [RFC-1123](#) requirements.

Note

If you need to change the name but the **Name** field is not editable, run the CIMC validation again to fix this issue.

- d) For **Type**, choose **Primary**.

The first nodes of the cluster must be set to **Primary**. You will add the secondary nodes in a later step if required for higher scale.

- e) In the **Data Network** area, enter the node's data network information.

Enter the data network IP address, netmask, and gateway. Optionally, you can also enter the VLAN ID for the network. Leave the VLAN ID field blank if your configuration does not require VLAN. If you chose **BGP** for **Cluster Connectivity**, enter the ASN.

If you enabled IPv6 functionality in a previous page, you must also enter the IPv6 address, netmask, and gateway.

Note

If you want to enter IPv6 information, you must do so during the cluster bootstrap process. To change the IP address configuration later, you would need to redeploy the cluster.

All nodes in the cluster must be configured with either only IPv4, only IPv6, or dual stack IPv4/IPv6.

- f) If you chose **BGP** for **Cluster Connectivity**, then in the **BGP peer details** area, enter the peer's IPv4 address and ASN.

You can click + **Add IPv4 BGP peer** to add additional peers.

If you enabled IPv6 functionality in a previous page, you must also enter the peer's IPv6 address and ASN.

g) Click **Save** to save the changes.

Step 21

In the **Node Details** screen, click **Add Node** to add the second node to the cluster.

If you are deploying a single-node cluster, skip this step.

Edit Node

General

Name *

nd-node1

Serial Number *

E5998163D6F0

Type *

Primary

Management Network ⓘ

IPv4 Address/Mask *

172.23.141.129/21

IPv4 Gateway *

172.23.136.1

IPv6 Address/Mask

IPv6 Gateway

Data Network ⓘ

IPv4 Address/Mask *

172.31.140.68/21

IPv4 Gateway *

172.31.136.1

IPv6 Address/Mask

IPv6 Gateway

VLAN ⓘ

Enable BGP ☐

Cancel

Save

- a) In the **Deployment Details** area, provide the **Management IP Address** and **Password** for the second node

You defined the management network information and the password during the initial node configuration steps.

- b) Click **Validate** to verify connectivity to the node.

The node's **Serial Number** and the **Management Network** information are automatically populated after connectivity is validated.

- c) Provide the **Name** for the node.
d) From the **Type** dropdown, select `Primary`.

The first 3 nodes of the cluster must be set to `Primary`. You will add the secondary nodes in a later step if required for higher scale.

- e) In the **Data Network** area, provide the node's **Data Network** information.

You must provide the data network IP address, netmask, and gateway. Optionally, you can also provide the VLAN ID for the network. For most deployments, you can leave the VLAN ID field blank.

If you had enabled IPv6 functionality in a previous screen, you must also provide the IPv6 address, netmask, and gateway.

Note

If you want to provide IPv6 information, you must do it during cluster bootstrap process. To change IP configuration later, you would need to redeploy the cluster.

All nodes in the cluster must be configured with either only IPv4, only IPv6, or dual stack IPv4/IPv6.

- f) (Optional) If your cluster is deployed in L3 mode, **Enable BGP** for the data network.

BGP configuration is required for the persistent IP addresses feature. This feature is described in more detail in [BGP configuration and persistent IP addresses, on page 46](#) and the "Persistent IP Addresses" sections of the [Cisco Nexus Dashboard User Guide](#).

Note

You can enable BGP at this time or in the Nexus Dashboard GUI after the cluster is deployed.

If you choose to enable BGP, you must also provide the following information:

- **ASN** (BGP Autonomous System Number) of this node.
You can configure the same ASN for all nodes or a different ASN per node.
- For IPv6-only, the **Router ID** of this node.
The router ID must be an IPv4 address, for example `1.1.1.1`
- **BGP Peer Details**, which includes the peer's IPv4 or IPv6 address and peer's ASN.

- g) Click **Save** to save the changes.
h) Repeat this step for the final (third) primary node of the cluster.

Step 22

(Optional) Repeat the previous step to enter information about any additional secondary or standby nodes.

Note

To support higher scale, you must provide a sufficient number of secondary nodes during deployment. Refer to the [Nexus Dashboard Cluster Sizing](#) tool for exact number of additional secondary nodes required for your specific use case.

You can choose to add the standby nodes now or at a later time after the cluster is deployed.

Step 23 In the **Node Details** page, verify the information that you entered, then click **Next**.

Step 24 In the **Persistent IPs** page, if you want to add more persistent IP addresses, click + **Add Data Service IP Address**, enter the IP address, and click the checkmark icon. Repeat this step as many times as desired, then click **Next**.

You must configure the minimum number of required persistent IP addresses during the bootstrap process. This step enables you to add more persistent IP addresses if desired.

Step 25 In the **Summary** page, review and verify the configuration information, click **Save**, and click **Continue** to confirm the correct deployment mode and proceed with building the cluster.

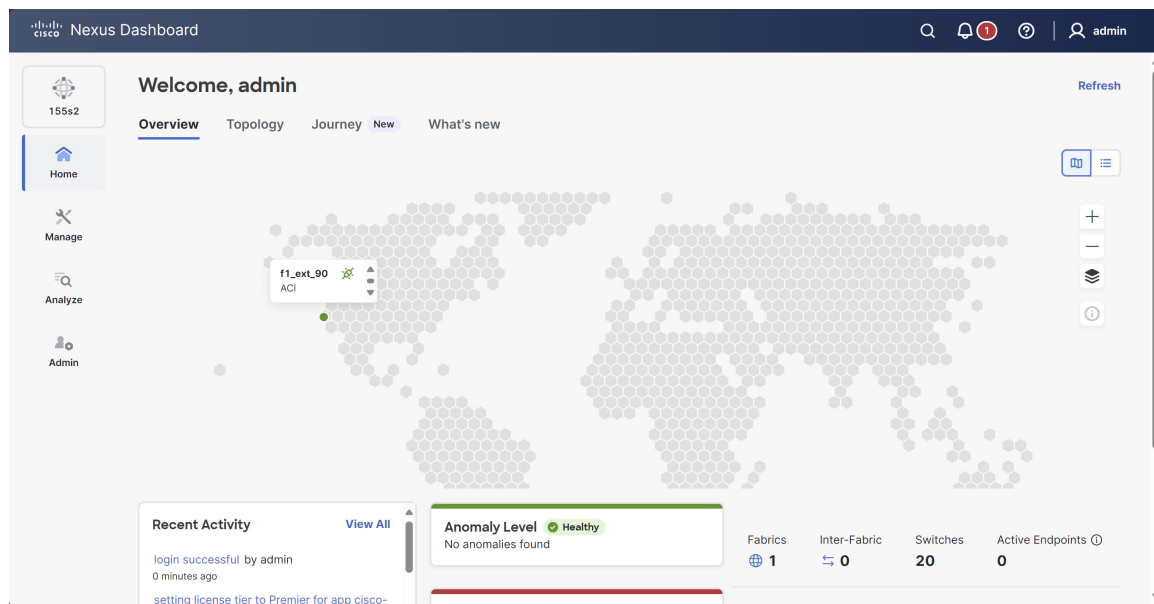
During the node bootstrap and cluster bring-up, the overall progress as well as each node's individual progress will be displayed in the UI. If you do not see the bootstrap progress advance, manually refresh the page in your browser to update the status.

It may take up to 60 minutes or more for the cluster to form, depending on the number of nodes in the cluster, and all the features to start. When cluster configuration is complete, the page will reload to the Nexus Dashboard GUI.

Step 26 Verify that the cluster is healthy.

After the cluster becomes available, you can access it by browsing to any one of your nodes' management IP addresses. The default password for the `admin` user is the same as the `rescue-user` password you chose for the first node. During this time, the UI will display a banner at the top stating "Service Installation is in progress, Nexus Dashboard configuration tasks are currently disabled".

After all the cluster is deployed and all services are started, you can look at the **Anomaly Level** on the **Home > Overview** page to ensure the cluster is healthy:



Alternatively, you can log in to any one node using SSH as the `rescue-user` using the password you entered during node deployment and using the `acs health` command to see the status:

- While the cluster is converging, you may see the following output:

```
$ acs health
k8s install is in-progress

$ acs health
k8s services not in desired state - [...]
```

```
$ acs health
k8s: Etcd cluster is not ready
```

- When the cluster is up and running, the following output will be displayed:

```
$ acs health
All components are healthy
```

Note

In some situations, you might power cycle a node (power it off and then back on) and find it stuck in this stage:

```
deploy base system services
```

This is due to an issue with `etcd` on the node after a reboot of the physical Nexus Dashboard cluster.

To resolve the issue, enter the `acs reboot clean` command on the affected node.

Step 27 (Optional) Connect your Cisco Nexus Dashboard cluster to Cisco Intersight for added visibility and benefits. Refer to [Working with Cisco Intersight](#) for detailed steps.

Step 28 After you have deployed Nexus Dashboard, see the [collections page](#) for this release for configuration information.

What to do next

The next task is to create the fabrics and fabric groups. See the *Creating Fabrics and Fabric Groups* article for this release on the [Cisco Nexus Dashboard collections page](#).

Deploy Nexus Dashboard Directly in VMware ESXi

This section describes how to deploy Cisco Nexus Dashboard cluster directly in VMware ESXi. If you prefer to deploy using vCenter, follow the steps described in [Deploy Nexus Dashboard Directly in VMware ESXi, on page 135](#) instead.

Before you begin

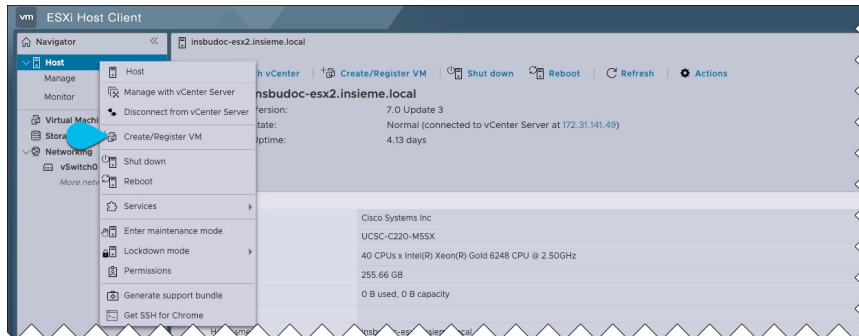
- Ensure that you meet the requirements and guidelines described in [Prerequisites and guidelines for deploying the Nexus Dashboard cluster in VMware ESX, on page 119](#).

Procedure

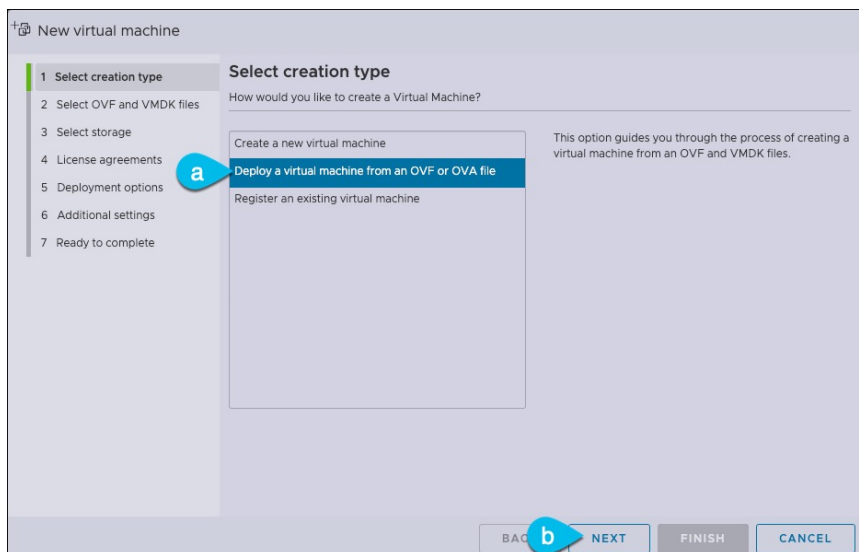
- Step 1** Obtain the Cisco Nexus Dashboard OVA image.
- Browse to the Software Download page.
<https://software.cisco.com/download/home/286327743/type/286328258/>
 - Choose the Nexus Dashboard release version you want to download.
 - Click the **Download** icon next to the Nexus Dashboard OVA image (`nd-dk9.<version>.ova`).
- Step 2** Log in to your VMware ESXi.

Depending on the version of your ESXi server, the location and order of configuration screens may differ slightly. The following steps provide deployment details using VMware ESXi 7.0.

Step 3 Right-click the host and select **Create/Register VM**.



Step 4 In the **Select creation type** screen, choose **Deploy a virtual machine from an OVF or OVA file**, then click **Next**.



Step 5 In the **Select OVF and VMDK files** screen, provide the virtual machine name (for example, `nd-ova-node1`) and the OVA image you downloaded in the first step, then click **Next**.

Step 6 In the **Select storage** screen, choose the datastore for the VM, then click **Next**.

Step 7 In the **Select OVF and VMDK files** screen, provide the virtual machine name (for example, `nd-node1`) and the OVA image you downloaded in the first step, then click **Next**.

Step 8 Specify the **Deployment options**.

In the **Deployment options** screen, provide the following:

- From the **Network mappings** dropdowns, choose the networks for the Nexus Dashboard management (`mgmt0`) and data (`fabric0`) interfaces.

Nexus Dashboard networks are described in [General prerequisites and guidelines, on page 11](#).

- From the **Deployment type** dropdown, choose the node profile (`App` or `Data`).

Node profiles are described in [Prerequisites and guidelines for deploying the Nexus Dashboard cluster in VMware ESX, on page 119](#).

- For **Disk provisioning** type, choose `Thick`.
- Disable the **Power on automatically** option.

Step 9 In the **Ready to complete** screen, verify that all information is accurate and click **Finish** to begin deploying the first node.

Step 10 Repeat previous steps to deploy the second and third nodes.

Note

If you are deploying a single-node cluster, you can skip this step.

You do not need to wait for the first node deployment to complete, you can begin deploying the other two nodes simultaneously.

Step 11 Wait for the VM(s) to finish deploying.

Step 12 Ensure that the VMware Tools periodic time synchronization is disabled, then start the VMs.

To disable time synchronization:

- Right-click the node's VM and select **Edit Settings**.
- In the **Edit Settings** window, select the **VM Options** tab.
- Expand the **VMware Tools** category and uncheck the **Synchronize guest time with host** option.

Step 13 Open one of the node's console and configure the node's basic information.

- Begin initial setup.

You will be prompted to run the first-time setup utility:

```
[ OK ] Started atomix-boot-setup.
      Starting Initial cloud-init job (pre-networking)...
      Starting logrotate...
      Starting logwatch...
      Starting keyhole...
[ OK ] Started keyhole.
[ OK ] Started logrotate.
[ OK ] Started logwatch.
```

Press any key to run first-boot setup on this console...

- Enter and confirm the `admin` password

This password will be used for the `rescue-user` SSH login as well as the initial GUI password.

Note

You must provide the same password for all nodes or the cluster creation will fail.

```
Admin Password:
Reenter Admin Password:
```

- Enter the management network information.

```
Management Network:
  IP Address/Mask: 192.168.9.172/24
  Gateway: 192.168.9.1
```

- For the first node only, designate it as the "Cluster Leader".

You will log into the cluster leader node to finish configuration and complete cluster creation.

```
Is this the cluster leader?: y
```

- e) Review and confirm the entered information.

You will be asked if you want to change the entered information. If all the fields are correct, choose **n** to proceed. If you want to change any of the entered information, enter **y** to re-start the basic configuration script.

```
Please review the config
Management network:
  Gateway: 192.168.9.1
  IP Address/Mask: 192.168.9.172/24
Cluster leader: no

Re-enter config? (y/N): n
```

Step 14 Repeat previous steps to deploy the additional nodes.

If you are deploying a single-node cluster, you can skip this step.

For multi-node clusters, you must deploy two additional **Primary** nodes and as many **Secondary** nodes as required by your specific use case. The total number of required nodes is available in the [Nexus Dashboard Capacity Planning](#) tool.

You do not need to wait for the first node configuration to complete, you can begin configuring the other two nodes simultaneously.

Note

You must provide the same password for all nodes or the cluster creation will fail.

The steps to deploy additional nodes are identical with the only exception being that you must indicate that they are not the **Cluster Leader**.

Step 15 Open your browser and navigate to `https://<node-mgmt-ip>` to open the GUI.

The rest of the configuration workflow takes place from one of the node's GUI. You can choose any one of the nodes you deployed to begin the bootstrap process and you do not need to log in to or configure the other two nodes directly.

Enter the password you entered in a previous step and click **Login**

Step 16 After logging in, review the **Meet Nexus Dashboard** welcome screens, then proceed to the **Journey** screen. Click on **Go** under the **Cluster bringup** step to begin the bootstrap process.

Step 17 Enter the requested information in the **Basic Information** page of the **Cluster Bringup** wizard.

- a) For **Cluster Name**, enter a name for this Nexus Dashboard cluster.

The cluster name must follow the [RFC-1123](#) requirements.

- b) For **Select the Nexus Dashboard Implementation type**, choose either **LAN** or **SAN** then click **Next**.

Step 18 Enter the requested information in the **Configuration** page of the **Cluster Bringup** wizard.

- (Optional) If you want to enable IPv6 functionality for the cluster, put a check in the **Enable IPv6** checkbox.
- Click **+Add DNS provider** to add one or more DNS servers, enter the DNS provider IP address, then click the checkmark icon.
- (Optional) Click **+Add DNS search domain** to add a search domain, enter the DNS search domain IP address, then click the checkmark icon.
- (Optional) If you want to enable NTP server authentication, put a check in the **NTP Authentication** checkbox.
- If you enabled NTP authentication, click **+ Add Key**, enter the required information, and click the checkmark icon to save the information.
 - **Key**—Enter the NTP authentication key, which is a cryptographic key that is used to authenticate the NTP traffic between the Nexus Dashboard and the NTP servers. You will define the NTP servers in the following step, and multiple NTP servers can use the same NTP authentication key.

- **ID**—Enter a key ID for the NTP host. Each NTP key must be assigned a unique key ID, which is used to identify the appropriate key to use when verifying the NTP packet.
- **Authentication Type**—Choose authentication type for the NTP key.
- Put a check in the **Trusted** checkbox if you want this key to be trusted. Untrusted keys cannot be used for NTP authentication.

For the complete list of NTP authentication requirements and guidelines, see [General prerequisites and guidelines, on page 11](#).

If you want to enter additional NTP keys, click + **Add Key** again and enter the information.

- f) If you enabled NTP authentication, click +**Add NTP Host Name/IP Address**, enter the required information, and click the checkmark icon to save the information.

- **NTP Host**—Enter an IP address; fully qualified domain names (FQDN) are not supported.

- **Key ID**—Enter the key ID of the NTP key you defined in the previous substep.

If NTP authentication is disabled, this field is grayed out.

- Put a check in the **Preferred** checkbox if you want this host to be preferred.

Note

If the node into which you are logged in is configured with only an IPv4 address, but you have checked **Enable IPv6** in a previous step and entered an IPv6 address for an NTP server, you will get the following validation error:

NTP Host*	Key ID	Preferred	
2001:420:28e:202a:5054:ff:fe6f:b3f6	true		 
 Add NTP Host Name/IP Address			

 Could not validate one or more hosts Can not reach NTP on Management Network

This is because the node does not have an IPv6 address yet and is unable to connect to an IPv6 address of the NTP server. You will enter IPv6 address in the next step. In this case, enter the other required information as described in the following steps and click **Next** to proceed to the next page where you will enter IPv6 addresses for the nodes.

If you want to enter additional NTP servers, click +**Add NTP Host Name/IP Address** again and enter the information.

- g) For **Proxy Server**, enter the URL or IP address of a proxy server.

For clusters that do not have direct connectivity to Cisco cloud, we recommend configuring a proxy server to establish the connectivity. This allows you to mitigate risk from exposure to non-conformant hardware and software in your fabrics.

You can click +**Add Ignore Host** to enter one or more destination IP addresses for which traffic will skip using the proxy.

The proxy server must permit these URLs:

```
svc.intersight.com
svc-static1.intersight.com
svc-static1.ucs-connect.com
```

If you do not want to configure a proxy, click **Skip Proxy** then click **Confirm**.

- h) (Optional) If your proxy server requires authentication, put a check in the **Authentication required for Proxy** checkbox and enter the login credentials.

- i) (Optional) Expand the **Advanced Settings** category and change the settings if required.

Under advanced settings, you can configure these settings:

- **App Network**—The address space used by the application's services running in the Nexus Dashboard. Enter the IP address and netmask.
- **Service Network**—An internal network used by Nexus Dashboard and its processes. Enter the IP address and netmask.
- **App Network IPv6**—If you put a check in the **Enable IPv6** checkbox earlier, enter the IPv6 subnet for the app network.
- **Service Network IPv6**—If you put a check in the **Enable IPv6** checkbox earlier, enter the IPv6 subnet for the service network.

For more information about the application and service networks, see [General prerequisites and guidelines, on page 11](#).

- j) Click **Next**.

Step 19

In the **Node Details** page, update the first node's information.

You have defined the Management network and IP address for the node into which you are currently logged in during the initial node configuration in earlier steps, but you must also enter the Data network information for the node before you can proceed with adding the other `primary` nodes and creating the cluster.

- a) For **Cluster Connectivity**, if your cluster is deployed in L3 mode, choose **BGP**. Otherwise, choose **L2**.

BGP configuration is required for the persistent IP addresses feature used by telemetry. This feature is described in more detail in the [BGP configuration and persistent IP addresses, on page 46](#) and [Nexus Dashboard persistent IP addresses, on page 40](#) sections.

Note

You can enable BGP at this time or in the Nexus Dashboard GUI after the cluster is deployed. All remaining nodes need to configure BGP if it is configured. You must enable BGP now if the data network of nodes have different subnets.

- b) Click the **Edit** button next to the first node.

The node's **Serial Number**, **Management Network** information, and **Type** are automatically populated, but you must enter the other information.

- c) For **Name**, enter a name for the node.

The node's **Name** will be set as its hostname, so it must follow the [RFC-1123](#) requirements.

Note

If you need to change the name but the **Name** field is not editable, run the CIMC validation again to fix this issue.

- d) For **Type**, choose **Primary**.

The first nodes of the cluster must be set to **Primary**. You will add the secondary nodes in a later step if required for higher scale.

- e) In the **Data Network** area, enter the node's data network information.

Enter the data network IP address, netmask, and gateway. Optionally, you can also enter the VLAN ID for the network. Leave the VLAN ID field blank if your configuration does not require VLAN. If you chose **BGP** for **Cluster Connectivity**, enter the ASN.

If you enabled IPv6 functionality in a previous page, you must also enter the IPv6 address, netmask, and gateway.

Note

If you want to enter IPv6 information, you must do so during the cluster bootstrap process. To change the IP address configuration later, you would need to redeploy the cluster.

All nodes in the cluster must be configured with either only IPv4, only IPv6, or dual stack IPv4/IPv6.

- f) If you chose **BGP** for **Cluster Connectivity**, then in the **BGP peer details** area, enter the peer's IPv4 address and ASN.

You can click + **Add IPv4 BGP peer** to add additional peers.

If you enabled IPv6 functionality in a previous page, you must also enter the peer's IPv6 address and ASN.

- g) Click **Save** to save the changes.

Step 20

In the **Node Details** screen, click **Add Node** to add the second node to the cluster.

If you are deploying a single-node cluster, skip this step.

Edit Node

General

Name *

nd-node1

Serial Number *

E5998163D6F0

Type *

Primary

Management Network ⓘ

IPv4 Address/Mask *

172.23.141.129/21

IPv4 Gateway *

172.23.136.1

IPv6 Address/Mask

IPv6 Gateway

Data Network ⓘ

IPv4 Address/Mask *

172.31.140.68/21

IPv4 Gateway *

172.31.136.1

IPv6 Address/Mask

IPv6 Gateway

VLAN ⓘ

Enable BGP ☐

Cancel

Save

- a) In the **Deployment Details** area, provide the **Management IP Address** and **Password** for the second node

You defined the management network information and the password during the initial node configuration steps.

- b) Click **Validate** to verify connectivity to the node.

The node's **Serial Number** and the **Management Network** information are automatically populated after connectivity is validated.

- c) Provide the **Name** for the node.
d) From the **Type** dropdown, select `Primary`.

The first 3 nodes of the cluster must be set to `Primary`. You will add the secondary nodes in a later step if required for higher scale.

- e) In the **Data Network** area, provide the node's **Data Network** information.

You must provide the data network IP address, netmask, and gateway. Optionally, you can also provide the VLAN ID for the network. For most deployments, you can leave the VLAN ID field blank.

If you had enabled IPv6 functionality in a previous screen, you must also provide the IPv6 address, netmask, and gateway.

Note

If you want to provide IPv6 information, you must do it during cluster bootstrap process. To change IP configuration later, you would need to redeploy the cluster.

All nodes in the cluster must be configured with either only IPv4, only IPv6, or dual stack IPv4/IPv6.

- f) (Optional) If your cluster is deployed in L3 mode, **Enable BGP** for the data network.

BGP configuration is required for the persistent IP addresses feature. This feature is described in more detail in [BGP configuration and persistent IP addresses, on page 46](#) and the "Persistent IP Addresses" sections of the [Cisco Nexus Dashboard User Guide](#).

Note

You can enable BGP at this time or in the Nexus Dashboard GUI after the cluster is deployed.

If you choose to enable BGP, you must also provide the following information:

- **ASN** (BGP Autonomous System Number) of this node.
You can configure the same ASN for all nodes or a different ASN per node.
- For IPv6-only, the **Router ID** of this node.
The router ID must be an IPv4 address, for example `1.1.1.1`
- **BGP Peer Details**, which includes the peer's IPv4 or IPv6 address and peer's ASN.

- g) Click **Save** to save the changes.
h) Repeat this step for the final (third) primary node of the cluster.

Step 21

(Optional) Repeat the previous step to enter information about any additional secondary or standby nodes.

Note

To support higher scale, you must provide a sufficient number of secondary nodes during deployment. Refer to the [Nexus Dashboard Cluster Sizing](#) tool for exact number of additional secondary nodes required for your specific use case.

You can choose to add the standby nodes now or at a later time after the cluster is deployed.

Step 22 In the **Node Details** page, verify the information that you entered, then click **Next**.

Step 23 In the **Persistent IPs** page, if you want to add more persistent IP addresses, click + **Add Data Service IP Address**, enter the IP address, and click the checkmark icon. Repeat this step as many times as desired, then click **Next**.

You must configure the minimum number of required persistent IP addresses during the bootstrap process. This step enables you to add more persistent IP addresses if desired.

Step 24 In the **Summary** page, review and verify the configuration information, click **Save**, and click **Continue** to confirm the correct deployment mode and proceed with building the cluster.

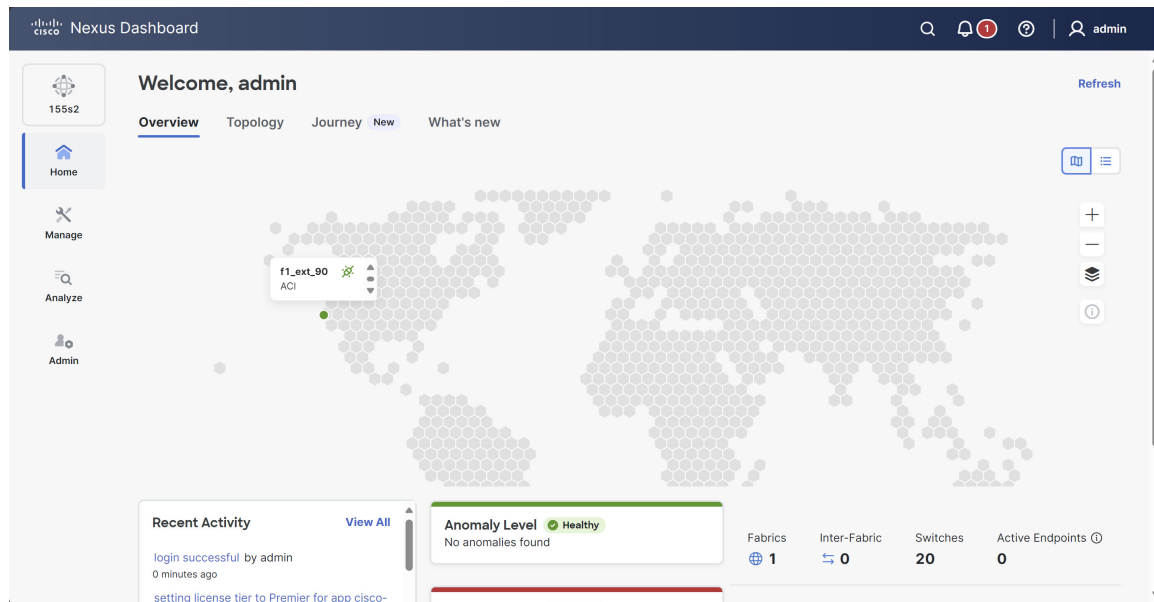
During the node bootstrap and cluster bring-up, the overall progress as well as each node's individual progress will be displayed in the UI. If you do not see the bootstrap progress advance, manually refresh the page in your browser to update the status.

It may take up to 60 minutes or more for the cluster to form, depending on the number of nodes in the cluster, and all the features to start. When cluster configuration is complete, the page will reload to the Nexus Dashboard GUI.

Step 25 Verify that the cluster is healthy.

After the cluster becomes available, you can access it by browsing to any one of your nodes' management IP addresses. The default password for the `admin` user is the same as the `rescue-user` password you chose for the first node. During this time, the UI will display a banner at the top stating "Service Installation is in progress, Nexus Dashboard configuration tasks are currently disabled".

After all the cluster is deployed and all services are started, you can look at the **Anomaly Level** on the **Home > Overview** page to ensure the cluster is healthy:



Alternatively, you can log in to any one node using SSH as the `rescue-user` using the password you entered during node deployment and using the `acs health` command to see the status:

- While the cluster is converging, you may see the following output:

```
$ acs health
k8s install is in-progress

$ acs health
k8s services not in desired state - [...]
```

```
$ acs health  
k8s: Etcd cluster is not ready
```

- When the cluster is up and running, the following output will be displayed:

```
$ acs health  
All components are healthy
```

Note

In some situations, you might power cycle a node (power it off and then back on) and find it stuck in this stage:

```
deploy base system services
```

This is due to an issue with `etcd` on the node after a reboot of the physical Nexus Dashboard cluster.

To resolve the issue, enter the `acs reboot clean` command on the affected node.

Step 26 (Optional) Connect your Cisco Nexus Dashboard cluster to Cisco Intersight for added visibility and benefits. Refer to [Working with Cisco Intersight](#) for detailed steps.

Step 27 After you have deployed Nexus Dashboard, see the [collections page](#) for this release for configuration information.



CHAPTER 7

Deploying in Linux KVM

- [Prerequisites and guidelines for deploying the Nexus Dashboard cluster in Linux KVM, on page 147](#)
- [About the deployment script for QCOW2 images, on page 149](#)
- [Deploy Nexus Dashboard in Linux KVM, on page 157](#)

Prerequisites and guidelines for deploying the Nexus Dashboard cluster in Linux KVM



Note Refer to the [Cisco Nexus Dashboard on KVM](#) document for platform-level guidance for maintaining Nexus Dashboard availability when users patch, upgrade, or otherwise modify the Linux KVM host.

Before you proceed with deploying the Nexus Dashboard cluster in a Linux KVM, the KVM must meet these prerequisites and you must follow these guidelines:

- The KVM form factor must support your scale requirements.
Scale support and co-hosting vary based on the cluster form factor. You can use the [Nexus Dashboard Capacity Planning](#) tool to verify that the virtual form factor satisfies your deployment requirements.
- Review and complete the general prerequisites described in [Prerequisites and Guidelines, on page 11](#).
- The CPU family used for the Nexus Dashboard VMs must support the AVX instruction set.
- The KVM must have enough system resources, and each node requires a dedicated disk partition. See [Understanding system resources, on page 149](#) for more information.
- The disk must have I/O latency of 20ms or less.
See [Verify the I/O latency of a Linux KVM storage device, on page 148](#)
- Cisco does not support the use of nested virtualization environments. Deploying Nexus Dashboard on a virtual machine that is itself running on a virtualized hypervisor (for example, KVM on ESXi) is not a supported configuration and may result in performance degradation or system instability.
- By default, many RAID controllers are set to `Write Through` mode. In this mode, the controller ensures that the data is physically written to the SSD flash before confirming success to the operating system.

Because of the "handshake" required for every synchronous write, disk I/O latencies can spike more than the recommended value. This usually happens when Telemetry is enabled on Nexus Dashboard, which requires a high disk I/O.

To optimize disk I/O performance, the RAID controller's onboard RAM must be used as a high-speed buffer. Make sure that the "Write Policy" is set to `Write Back (WB)`. This allows the controller to acknowledge writes immediately using its onboard cache.

- KVM deployments are supported for NX-OS and ACI fabrics, as well as for SAN deployments.
- You must deploy in Red Hat Enterprise Linux 8.8, 8.10, or 9.4. In addition, Cisco strongly recommends that all nodes are running on the same Linux version.
- In order for Nexus Dashboard to be up and running on OS reboot scenarios, you must add the UUIDs in the `fstab` `conf` files of your RHEL host operating system, which is the only way to preserve the Nexus Dashboard upon a reboot of the RHEL operating system.
- You must also configure the following required network bridges at the host level for Nexus Dashboard deployments:
 - Management Network Bridge (mgmt-bridge): The external network to manage Nexus Dashboard.
 - Data Network Bridge (data-bridge): The internal network used to form clustering within Nexus Dashboard.
- We recommend that each Nexus Dashboard node is deployed in a different KVM hypervisor.

Verify the I/O latency of a Linux KVM storage device

When you deploy a Nexus Dashboard cluster in a Linux KVM, the storage device of the KVM must have a latency under 20ms. This is done in the Performance Validation step if you use the `kvm_deployer` deployment script for QCOW2 images, as described in [About the deployment script for QCOW2 images, on page 149](#).

Follow these steps if you want to manually verify the I/O latency of a Linux KVM storage device.

Procedure

- | | |
|---------------|--|
| Step 1 | Create a test directory.
For example, create a directory named <code>test-data</code> . |
| Step 2 | Run the Flexible I/O tester (FIO).

<pre># fio --rw=write --ioengine=sync --fdatasync=1 --directory=test-data --size=22m --bs=2300 --name=mytest</pre> |
| Step 3 | After you use the command, confirm that the <code>99.00th=[value]</code> in the <code>fsync/fdatasync/sync_file_range</code> section is under 20ms. |

Understanding system resources

When deploying a Nexus Dashboard cluster in Linux KVM, the KVM must have enough system resources. There are multiple form factors supported with a virtual Nexus Dashboard KVM, and the amount of system resources needed for each node differs based on the form factor.

Table 26: Per node resource requirements

Form factor	Number of vCPUs	RAM size	Disk size
1-node KVM (app)	16	64 GB	550 GB
1-node KVM (data)	32	128 GB	3 TB
3-node KVM (app)	16	64 GB	550 GB
3-node KVM (data)	32	128 GB	3 TB

You will need to know the information above for your form factor when you go through the procedures in [Deploy Nexus Dashboard in Linux KVM, on page 157](#).

About the deployment script for QCOW2 images

Nexus Dashboard supports KVM virtualization across various environments, including RHEL and Nutanix. Nexus Dashboard deployments are typically done using QCOW2 images; however, the diversity of KVM variants and host configurations creates challenges in tracking supported versions and ensuring optimal device emulation for performance.

The deployment script for QCOW2 images addresses these challenges by providing a deployment validation and orchestration script that is packaged with the QCOW2 image. This script is designed to streamline the deployment process, ensure compatibility, and help maintain consistent performance standards across supported platforms.

Key Features

These are the key features for the deployment script for QCOW2 images:

- **Platform Compatibility Check:** Verifies that the underlying host operating system is a supported distribution and version.
- **Resource Validation:** Assesses whether the host meets minimum compute requirements, including CPU, memory, and storage resources.
- **Performance Validation:**
 - **Disk I/O and Network Latency:** Measures key performance indicators such as disk I/O and network I/O latency, ensuring they fall within qualified ranges.
 - **Device Emulation:** Confirms that the correct device emulation (such as network interfaces and virtual disk devices) is available for optimal performance.

- **Deployment Orchestration:** Automates the deployment process for both single-node and multi-node clusters, orchestrating virtual machines with the appropriate device mappings (such as network and storage devices).

Because storage is the most common bottleneck, the script categorizes storage into "Performance Tiers":

- **Disk Type & Placement (SSD/NVMe):** Validates for high-IOPs workloads.
- **RAID vs. JBOD:** Validates write speed.
- **Remote Storage (iSCSI/FC/NFS):** Performs a "Pre-Flight Latency Check":
 - **Sync Latency:** Must be < 10ms for standard apps and < 1ms for high-performance databases.
 - **IOPS Validation:** Uses a temporary `fio` test to ensure the required IOS.

Prerequisites for the deployment script for QCOW2 images

These are the prerequisites before using the deployment script for QCOW2 images.

Requirement	Details
ND qcow2 image	A local path to the Nexus Dashboard .qcow2 image file
Network info	3 management IPs + a gateway
Go toolchain	Go 1.22+ to build from source (or use a pre-built binary)

Use the deployment script for QCOW2 images

Procedure

Step 1 Locate the tar file with the deployment script.

The tar file should be available at the same location as the QCOW2 image. For example, in the [Software Download site for Nexus Dashboard](#) for the 4.3.1 release and later, you should see the tar file with the deployment script for QCOW2 images at the same area as the QCOW2 image itself, such as this:

```
nd-dk9.4.3.1.175.qcow2.tar.gz
```

Step 2 Download the tar file with the deployment script.

Step 3 Extract the tar file.

For example:

```
tar -xzf nd-dk9.4.3.1.175.qcow2.tar.gz
```

You should see output similar to the following after extracting the tar file and listing the contents in the directory:

```
#ls -lrt
total 24730576
-rwxr-xr-x 1 root root 35991736 May 5 23:35 kvm_deployer
```

```
-rw-r--r-- 1 root root 12651498496 May 17 03:44 nd-dk9.4.3.1.175.qcow2
-rw-r--r-- 1 root root 12636610743 May 18 22:44 nd-dk9.4.3.1.175.qcow2.tar.gz
```

Step 4 Use the **kvm_deployer config** command to create a yaml file to use with the deployment script.

Make the appropriate choices at each stage based on your deployment type.

- For example, to generate a one-node yaml file:

```
# ./kvm_deployer config

Configure Nexus Dashboard deployment:
Deployment Platform (kvm/nutanix): kvm
Image Path (qcow2): /home/nd-base/nd-dk9.4.3.1.58.qcow2
Cluster Size: 1
VM Type (app/data/large/xlarge): app
Admin Password:
Configure current node:
KVM Host:
Management Network Bridge: nd_mgmt
Data Network Bridge: nd_data
Disk Path (device or mount point, absolute path): /home/nd-node1
Nexus Dashboard:
VM Name (lowercase, numbers, dashes, max 63 chars): node1
Management Network (e.g., 192.168.1.100/24): 192.168.1.128/25
Management Gateway IP: 192.168.1.129
[ ] Configuration Summary:
Deployment Platform      : kvm
Image Path (qcow2)      : /home/nd-base/nd-dk9.4.3.1.58.qcow2
Admin Password          : [hidden]
VM Type                 : app
VMs:
  VM Name                : node1
  Disk Path              : /home/nd-node1
  Management Network Bridge : nd_mgmt
  Data Network Bridge    : nd_data
  Management Network     : 192.168.1.128/25
  Management Gateway IP  : 192.168.1.129
Proceed with current config? (y/n): y
File name to save configuration: kvm_1node_deploy.yaml
```

- To generate a three-node yaml file:

```
# ./kvm_deployer config

Configure Nexus Dashboard deployment:
Deployment Platform (kvm/nutanix): kvm
Image Path (qcow2): /home/nd-base/nd-dk9.4.3.1.175.qcow2
Cluster Size: 3
VM Type (app/data/large/xlarge): app
Admin Password:
Configure current node:
KVM Host:
Management Network Bridge: nd_mgmt
Data Network Bridge: nd_data
Disk Path (device or mount point, absolute path): /home/nd-node1
Nexus Dashboard:
VM Name (lowercase, numbers, dashes, max 63 chars): node1
Management Network (e.g., 192.168.1.100/24): 192.168.1.128/25
Management Gateway IP: 192.168.1.129
Configure remote to deploy vm (2):
KVM Host:
IP Address: 192.168.1.133
Username (ssh): root
. Password (ssh: root):
```

```

Management Network Bridge: nd_mgmt
Data Network Bridge: nd_data
Disk Path (device or mount point, absolute path): /home/nd-node2
Nexus Dashboard:
  VM Name (lowercase, numbers, dashes, max 63 chars): node2
  Management Network (e.g., 192.168.1.100/24): 192.168.1.128/25
  Management Gateway IP: 192.168.1.130
Configure remote to deploy vm (3):
KVM Host:
  IP Address: 192.168.1.133
  Username (ssh): root
  Password (ssh: root):
Management Network Bridge: nd_mgmt
Data Network Bridge: nd_data
Disk Path (device or mount point, absolute path): /home/nd-node3
Nexus Dashboard:
  VM Name (lowercase, numbers, dashes, max 63 chars): node3
  Management Network (e.g., 192.168.1.100/24): 192.168.1.128/25
  Management Gateway IP: 192.168.1.131
[ ] Configuration Summary:
Deployment Platform      : kvm
Image Path (qcow2)      : /home/nd-base/nd-dk9.4.3.1.175.qcow2
Admin Password          : [hidden]
VM Type                 : app
VMs:
  VM Name                : node1
  Disk Path              : /home/nd-node1
  Management Network Bridge : nd_mgmt
  Data Network Bridge    : nd_data
  Management Network      : 192.168.1.128/25
  Management Gateway IP   : 192.168.1.129
  Host OS IP Address      : 192.168.1.133
  Host OS Username        : root
  Host OS Password        : [hidden]
  VM Name                : node2
  Disk Path              : /home/nd-node2
  Management Network Bridge : nd_mgmt
  Data Network Bridge    : nd_data
  Management Network      : 192.168.1.128/25
  Management Gateway IP   : 192.168.1.130
  Host OS IP Address      : 192.168.1.133
  Host OS Username        : root
  Host OS Password        : [hidden]
  VM Name                : node3
  Disk Path              : /home/nd-node3
  Management Network Bridge : nd_mgmt
  Data Network Bridge    : nd_data
  Management Network      : 192.168.1.128/25
  Management Gateway IP   : 192.168.1.131
Proceed with current config? (y/n): y
File name to save configuration: kvm_3node_deploy.yaml

```

Step 5 (Optional) Use a text editor to verify the information saved in the yaml file, if necessary.

For example, for the three-node example above, the yaml file would have information similar to the following:

```

deployment_type: kvm
image_path: /home/nd-base/nd-dk9.4.3.1.175.qcow2
cluster_size: "3"
vms:
- name: node1
  disk_path: /home/nd-node1
  management_network: 192.168.1.128/25

```

```

management_gateway: 192.168.1.129
management_bridge: nd_mgmt
data_bridge: nd_data
host_ip: ""
host_user: ""
host_password: ""
- name: node2
  disk_path: /home/nd-node2
  management_network: 192.168.1.135/25
  management_gateway: 192.168.1.129
  management_bridge: nd_mgmt
  data_bridge: nd_data
  host_ip: 192.168.1.133
  host_user: root
  host_password: <host_password>
- name: node3
  disk_path: /home/nd-node3
  management_network: 192.168.1.136/25
  management_gateway: 192.168.1.129
  management_bridge: nd_mgmt
  data_bridge: nd_data
  host_ip: 192.168.1.133
  host_user: root
  host_password: <hostpword>
flavor: app
admin_password: <adminpword>

```

Step 6 Run the deployment script, using the yaml file as the input.

```
./kvm_deployer deploy --config kvm_3node_deploy.yaml
```

Output similar to the following displays:

```

Deploying VM node1 using kvm
Cleanup old images
[ ] Cleaned image /home/nd-nodel1/boot.img
[ ] Cleaned image /home/nd-nodel1/disk.img
[ ] Cleaned image /home/nd-nodel1/config-drive.img
Creating boot disk, it can take several minutes, pls wait
[ ] Boot disk created successfully
Creating raw disk, it can take several minutes, pls wait
[ ] Raw disk created successfully
Creating configuration drive
[ ] Configuration drive created successfully

Starting VM
[ ] VM started successfully
[ ] KVM VM: nodel1 deployed successfully
Deploying VM node2 using kvm
Copying file /home/nd-base/nd-dk9.4.3.1.175.qcow2 to /home/nd-node2/nd-dk9.4.3.1.175.qcow2
[ ] File /home/nd-base/nd-dk9.4.3.1.175.qcow2 copied to /home/nd-node2/nd-dk9.4.3.1.175.qcow2
Cleanup old images
[ ] Cleaned image /home/nd-node2/boot.img
[ ] Cleaned image /home/nd-node2/disk.img
[ ] Cleaned image /home/nd-node2/config-drive.img
Creating boot disk, it can take several minutes, pls wait
[ ] Boot disk created successfully
Creating raw disk, it can take several minutes, pls wait
[ ] Raw disk created successfully
Creating configuration drive
Copying file /home/nd-nodel1/node2-bootstrap.json to /home/nd-node2/apic-sn-cfg-mount/bootstrap.json
[ ] File /home/nd-nodel1/node2-bootstrap.json copied to /home/nd-node2/apic-sn-cfg-mount/bootstrap.json

[ ] Configuration drive created successfully

```

```

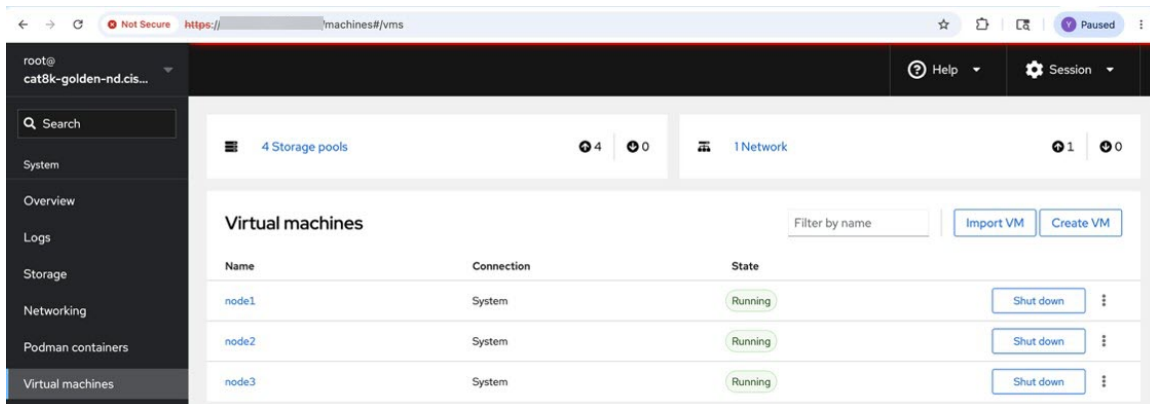
Starting VM
[ ] VM started successfully
[ ] KVM VM: node2 deployed successfully
Deploying VM node3 using kvm
Copying file /home/nd-base/nd-dk9.4.3.1.175.qcow2 to /home/nd-node3/nd-dk9.4.3.1.175.qcow2
[ ] File /home/nd-base/nd-dk9.4.3.1.175.qcow2 copied to /home/nd-node3/nd-dk9.4.3.1.175.qcow2
Cleanup old images
[ ] Cleaned image /home/nd-node3/boot.img
[ ] Cleaned image /home/nd-node3/disk.img
[ ] Cleaned image /home/nd-node3/config-drive.img
Creating boot disk, it can take several minutes, pls wait
[ ] Boot disk created successfully
Creating raw disk, it can take several minutes, pls wait
[ ] Raw disk created successfully
Creating configuration drive
Copying file /home/nd-node1/node3-bootstrap.json to /home/nd-node3/apic-sn-cfg-mount/bootstrap.json
[ ] File /home/nd-node1/node3-bootstrap.json copied to /home/nd-node3/apic-sn-cfg-mount/bootstrap.json

[ ] Configuration drive created successfully
Starting VM
[ ] VM started successfully
[ ] KVM VM: node3 deployed successfully

```

Step 7 Complete the post-deployment steps.

- a) **Verify VMs in RHEL host**— Open https://<RHEL_host_ip>:9090 and confirm all three VMs are powered on and healthy.



- b) **Cluster formation**— Once all the nodes are running, complete the Nexus Dashboard cluster setup through the Nexus Dashboard UI using the management IP addresses from your config.

← → ↺ Not Secure https:// 1/cluster-bringup ☆ 📄 📄 ⌵ Paused ⓘ

cluster

cisco

Nexus Dashboard

admin

⌵

Admin

← Journey

Cluster Bringup

1 Basic information

2 Configuration

3 Node details

4 Persistent IPs

5 Summary

Basic information

Provide basic information about this cluster

Cluster name cannot be changed after Cluster Bringup.

Changing this will require a cluster rebuild.

Cluster name *

Nd-KVM-133

Select the Nexus Dashboard implementation type

☒ LAN

Includes Cisco NX-OS, ACI, IOS XE, IOS XR and other 3rd party use cases

☐ SAN

Includes fiber channel and storage use cases

Cancel

Next

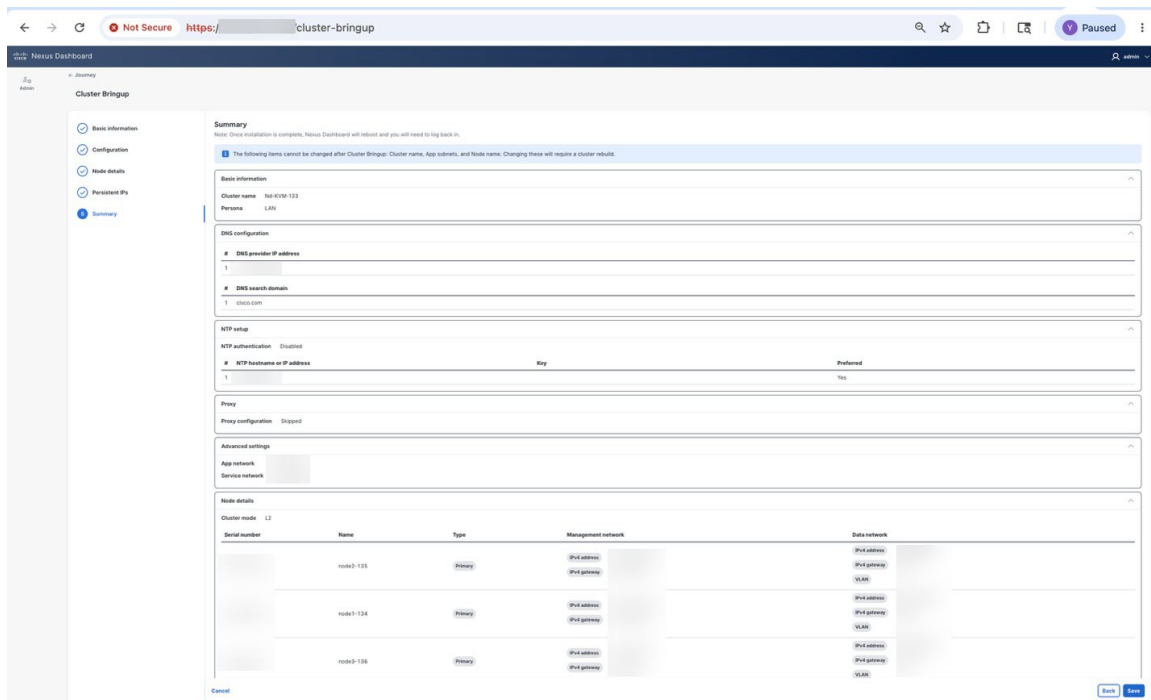
Cluster connectivity
* ⓘ

☒ L2
☐ BGP

Serial number	Name	Type	Management network	Data network	Configuration status ⓘ	
873DF24266F2	node1-134	Primary	IPv4 address: 10.10.10.10 IPv4 gateway: 10.10.10.1	IPv4 address: 10.10.10.10 IPv4 gateway: 10.10.10.1 VLAN: -	✓ Successful	
4EE78A04C3CE	node2-135	Primary	IPv4 address: 10.10.10.10 IPv4 gateway: 10.10.10.1	IPv4 address: 10.10.10.10 IPv4 gateway: 10.10.10.1 VLAN: -	✓ Successful	
WVT2922012Y	node3-136	Primary	IPv4 address: 10.10.10.10 IPv4 gateway: 10.10.10.1	IPv4 address: 10.10.10.10 IPv4 gateway: 10.10.10.1 VLAN: -	✓ Successful	

+ Add node

Cancel
Back
Validate



Optional post-deployment tasks

Troubleshooting

Symptom	Check
Authentication failure	Verify <code>username</code> and <code>password</code> in the config
VM not getting IP	Check that <code>cloud-init management_network</code> uses valid CIDR notation and the gateway is correct

Destroy the deployment

If you want to destroy the deployment, enter this command:

```
./kvm_deployer destroy --config <yaml-filename>.yaml
```

For example:

```
$ kvm_deployer destroy --config kvm_3node_deploy.yaml
Destroying VM node1 using nutanix
Destroying VM node1 using nutanix
Looking for VM: node1
Found VM node1 with UUID: 94fe9518-5196-4807-a6cc-183d1abe760e
Powering off VM: 94fe9518-5196-4807-a6cc-183d1abe760e
VM power off request submitted
Deleting VM: 94fe9518-5196-4807-a6cc-183d1abe760e
VM deletion task submitted: 8178938b-5014-474b-b168-34c3bf5b1577
Waiting for task completion: 8178938b-5014-474b-b168-34c3bf5b1577
Task status: RUNNING, progress: 0%
```

```

Task status: SUCCEEDED, progress: 100%
[ ] Task completed successfully (status: SUCCEEDED)
[ ] VM deleted successfully
[ ] VM node1 destroyed successfully
Destroying VM node2 using nutanix
Destroying VM node2 using nutanix
Looking for VM: node2
Found VM node2 with UUID: f1424251-182d-4518-9acc-f06ba89d4883
Powering off VM: f1424251-182d-4518-9acc-f06ba89d4883
VM power off request submitted
Deleting VM: f1424251-182d-4518-9acc-f06ba89d4883
VM deletion task submitted: a8d41687-3d84-4a38-91cc-5dfee0688351
Waiting for task completion: a8d41687-3d84-4a38-91cc-5dfee0688351
Task status: RUNNING, progress: 0%
Task status: SUCCEEDED, progress: 100%
[ ] Task completed successfully (status: SUCCEEDED)
[ ] VM deleted successfully
[ ] VM node2 destroyed successfully
Destroying VM node3 using nutanix
Destroying VM node3 using nutanix
Looking for VM: node3
Found VM node3 with UUID: 8cdc2037-fc83-43af-90f6-869c00aa2fdc
Powering off VM: 8cdc2037-fc83-43af-90f6-869c00aa2fdc
VM power off request submitted
Deleting VM: 8cdc2037-fc83-43af-90f6-869c00aa2fdc
VM deletion task submitted: 710c58a7-0728-4c4e-b63a-f6d90e3b2669
Waiting for task completion: 710c58a7-0728-4c4e-b63a-f6d90e3b2669
Task status: RUNNING, progress: 0%
Task status: SUCCEEDED, progress: 100%
[ ] Task completed successfully (status: SUCCEEDED)
[ ] VM deleted successfully
[ ] VM node3 destroyed successfully
$

```

Deploy Nexus Dashboard in Linux KVM

This section describes how to deploy Cisco Nexus Dashboard cluster in Linux KVM.

Before you begin

- Ensure that you meet the requirements and guidelines described in [Prerequisites and guidelines for deploying the Nexus Dashboard cluster in Linux KVM](#), on page 147.

Procedure

Step 1

Download the Cisco Nexus Dashboard image.

- Browse to the Software Download page.

<https://software.cisco.com/download/home/286327743/type/286328258>

- Click **Nexus Dashboard Software**.
- From the left sidebar, choose the Nexus Dashboard version you want to download.
- Download the Cisco Nexus Dashboard image for Linux KVM (nd-dk9.<version>.qcow2).

Step 2

Copy the image to the Linux KVM servers where you will host the nodes.

You can use `scp` to copy the image, for example:

```
# scp nd-dk9.<version>.qcow2 root@<kvm-host-ip>:/home/nd-base
```

The following steps assume you copied the image into the `/home/nd-base` directory.

Step 3

Make the following configurations on each KVM host:

- a) Edit `/etc/libvirt/qemu.conf` and make sure the user and group is correctly configured based on the ownership of the storage that you plan to use for the Nexus Dashboard deployment.

This is only required if you plan to use disk storage paths that are different from the default `libvirtd`.

- b) Edit `/etc/libvirt/libvirt.conf` and uncomment `uri_default`.
- c) Restart the `libvirtd` service after updating the configuration using the `systemctl restart libvirtd` command from root.

Step 4

Log in to your KVM host as the `root` user and perform the following steps to create the required disk images on each node.

As mentioned in [Understanding system resources, on page 149](#), you will need a total of 550 GB or 3 TB of SSD storage to create two disk images:

- Boot disk based on QCOW2 image that you downloaded
 - Data disk
- a) Verify that you have a directory with enough space to store the VM disks (for example, `/home/nd-node1`) or mount the storage disk (raw disk or LVM) to directory `/opt/cisco/nd`.
 - b) Create the following script as `/root/create_vm.sh` under the root directory.

Note

If you manually type this information, verify that there are no empty spaces present after any of these lines.

Create the script based on the information provided in [Understanding system resources, on page 149](#) for your form factor.

- For 1-node or 3-node KVM (app) form factors:

```
#!/bin/bash -ex

# Configuration
# Name of Nexus Dashboard Virtual machine
name=nd1

# Path of Nexus Dashboard QCOW2 image.
nd_qcow2=/home/nd-base/nd-dk9.4.2.1g.qcow2

# Disk Path to storage Boot and Data Disks.
data_disk=/opt/cisco/nd/data

# Management Network Bridge
mgmt_bridge=mgmt-bridge

# Data Network bridge
data_bridge=data-bridge

# Data Disk Size
data_size=500G

# CPU Cores
cpus=16
```

```
# Memory in units of MB.
memory=65536

# actual script
rm -rf $data_disk/boot.img
/usr/bin/qemu-img convert -f qcow2 -O raw $nd_qcow2 $data_disk/boot.img
rm -rf $data_disk/disk.img
/usr/bin/qemu-img create -f raw $data_disk/disk.img $data_size
virt-install \
--import \
--name $name \
--memory $memory \
--vcpus $cpus \
--os-type generic \
--osinfo detect=on,require=off \
--check path_in_use=off \
--disk path=${data_disk}/boot.img,format=raw,bus=virtio \
--disk path=${data_disk}/disk.img,format=raw,bus=virtio \
--network bridge=$mgmt_bridge,model=virtio \
--network bridge=$data_bridge,model=virtio \
--console pty,target_type=serial \
--noautoconsole \
--autostart
```

- For 1-node or 3-node KVM (**data**) form factors:

```
#!/bin/bash -ex

# Configuration
# Name of Nexus Dashboard Virtual machine
name=nd1

# Path of Nexus Dashboard QCOW2 image.
nd_qcow2=/home/nd-base/nd-dk9.4.2.1g.qcow2

# Disk Path to storage Boot and Data Disks.
data_disk=/opt/cisco/nd/data

# Management Network Bridge
mgmt_bridge=mgmt-bridge

# Data Network bridge
data_bridge=data-bridge

# Data Disk Size
data_size=3072G

# CPU Cores
cpus=32

# Memory in units of MB.
memory=131072

# actual script
rm -rf $data_disk/boot.img
/usr/bin/qemu-img convert -f qcow2 -O raw $nd_qcow2 $data_disk/boot.img
rm -rf $data_disk/disk.img
/usr/bin/qemu-img create -f raw $data_disk/disk.img $data_size
virt-install \
--import \
--name $name \
--memory $memory \
--vcpus $cpus \
--os-type generic \
```

```
--osinfo detect=on,require=off \
--check_path_in_use=off \
--disk path=${data_disk}/boot.img,format=raw,bus=virtio \
--disk path=${data_disk}/disk.img,format=raw,bus=virtio \
--network bridge=$mgmt_bridge,model=virtio \
--network bridge=$data_bridge,model=virtio \
--console pty,target_type=serial \
--noautoconsole \
--autostart
```

Note

If the automatic PCI assignment does not place the management and data interfaces at 00:03.0 and 00:04.0 respectively, connectivity might fail. In that case, replace these two `--network` lines:

```
--network bridge=$mgmt_bridge,model=virtio \
--network bridge=$data_bridge,model=virtio \
```

with these lines:

```
--network bridge=$mgmt_bridge,model=virtio,address.type=pci,address.bus=0,address.slot=3 \
--network bridge=$data_bridge,model=virtio,address.type=pci,address.bus=0,address.slot=4 \
```

Step 5 Make the `create_vm.sh` script executable and run it using these commands.

```
# chmod +x /root/create_vm.sh
# /root/create_vm.sh
```

Step 6 Repeat previous steps to deploy the second and third nodes, then start all VMs.

Note

If you are deploying a single-node cluster, you can skip this step.

Step 7 Open one of the node's console and configure the node's basic information.

a) Press any key to begin initial setup.

You will be prompted to run the first-time setup utility:

```
[ OK ] Started atomix-boot-setup.
      Starting Initial cloud-init job (pre-networking)...
      Starting logrotate...
      Starting logwatch...
      Starting keyhole...
[ OK ] Started keyhole.
[ OK ] Started logrotate.
[ OK ] Started logwatch.
```

Press any key to run first-boot setup on this console...

b) Enter and confirm the `admin` password

This password will be used for the `rescue-user` SSH login as well as the initial GUI password.

Note

You must provide the same password for all nodes or the cluster creation will fail.

```
Admin Password:
Reenter Admin Password:
```

c) Enter the management network information.

```
Management Network:
  IP Address/Mask: 192.168.9.172/24
  Gateway: 192.168.9.1
```

- d) For the first node only, designate it as the "Cluster Leader".

You will log into the cluster leader node to finish configuration and complete cluster creation.

```
Is this the cluster leader?: y
```

- e) Review and confirm the entered information.

You will be asked if you want to change the entered information. If all the fields are correct, choose `n` to proceed. If you want to change any of the entered information, enter `y` to re-start the basic configuration script.

```
Please review the config
Management network:
  Gateway: 192.168.9.1
  IP Address/Mask: 192.168.9.172/24
Cluster leader: yes

Re-enter config? (y/N): n
```

Step 8 Repeat previous step to configure the initial information for the second and third nodes.

You do not need to wait for the first node configuration to complete, you can begin configuring the other two nodes simultaneously.

Note

You must provide the same password for all nodes or the cluster creation will fail.

The steps to deploy the second and third nodes are identical with the only exception being that you must indicate that they are not the **Cluster Leader**.

Step 9 Wait for the initial bootstrap process to complete on all nodes.

After you provide and confirm management network information, the initial setup on the first node (`Cluster Leader`) configures the networking and brings up the UI, which you will use to add two other nodes and complete the cluster deployment.

```
Please wait for system to boot: [#####] 100%
System up, please wait for UI to be online.
```

```
System UI online, please login to https://192.168.9.172 to continue.
```

Step 10 Open your browser and navigate to `https://<node-mgmt-ip>` to open the GUI.

The rest of the configuration workflow takes place from one of the node's GUI. You can choose any one of the nodes you deployed to begin the bootstrap process and you do not need to log in to or configure the other two nodes directly.

Enter the password you entered in a previous step and click **Login**

Step 11 After logging in, review the **Meet Nexus Dashboard** welcome screens, then proceed to the **Journey** screen. Click on **Go** under the **Cluster bringup** step to begin the bootstrap process.

Step 12 Enter the requested information in the **Basic Information** page of the **Cluster Bringup** wizard.

- a) For **Cluster Name**, enter a name for this Nexus Dashboard cluster.

The cluster name must follow the [RFC-1123](#) requirements.

- b) For **Select the Nexus Dashboard Implementation type**, choose either **LAN** or **SAN** then click **Next**.

Step 13 Enter the requested information in the **Configuration** page of the **Cluster Bringup** wizard.

- a) (Optional) If you want to enable IPv6 functionality for the cluster, put a check in the **Enable IPv6** checkbox.
- b) Click **+Add DNS provider** to add one or more DNS servers, enter the DNS provider IP address, then click the checkmark icon.
- c) (Optional) Click **+Add DNS search domain** to add a search domain, enter the DNS search domain IP address, then click the checkmark icon.
- d) (Optional) If you want to enable NTP server authentication, put a check in the **NTP Authentication** checkbox.
- e) If you enabled NTP authentication, click **+ Add Key**, enter the required information, and click the checkmark icon to save the information.
 - **Key**—Enter the NTP authentication key, which is a cryptographic key that is used to authenticate the NTP traffic between the Nexus Dashboard and the NTP servers. You will define the NTP servers in the following step, and multiple NTP servers can use the same NTP authentication key.
 - **ID**—Enter a key ID for the NTP host. Each NTP key must be assigned a unique key ID, which is used to identify the appropriate key to use when verifying the NTP packet.
 - **Authentication Type**—Choose authentication type for the NTP key.
 - Put a check in the **Trusted** checkbox if you want this key to be trusted. Untrusted keys cannot be used for NTP authentication.

For the complete list of NTP authentication requirements and guidelines, see [General prerequisites and guidelines, on page 11](#).

If you want to enter additional NTP keys, click **+ Add Key** again and enter the information.

- f) If you enabled NTP authentication, click **+Add NTP Host Name/IP Address**, enter the required information, and click the checkmark icon to save the information.
 - **NTP Host**—Enter an IP address; fully qualified domain names (FQDN) are not supported.
 - **Key ID**—Enter the key ID of the NTP key you defined in the previous substep.
If NTP authentication is disabled, this field is grayed out.
 - Put a check in the **Preferred** checkbox if you want this host to be preferred.

Note

If the node into which you are logged in is configured with only an IPv4 address, but you have checked **Enable IPv6** in a previous step and entered an IPv6 address for an NTP server, you will get the following validation error:

NTP Host*	Key ID	Preferred	
2001:420:28e:202a:5054:ff:fe6f:b3f6		true	 
 Add NTP Host Name/IP Address			

 Could not validate one or more hosts Can not reach NTP on Management Network

This is because the node does not have an IPv6 address yet and is unable to connect to an IPv6 address of the NTP server. You will enter IPv6 address in the next step. In this case, enter the other required information as described in the following steps and click **Next** to proceed to the next page where you will enter IPv6 addresses for the nodes.

If you want to enter additional NTP servers, click **+Add NTP Host Name/IP Address** again and enter the information.

- g) For **Proxy Server**, enter the URL or IP address of a proxy server.

For clusters that do not have direct connectivity to Cisco cloud, we recommend configuring a proxy server to establish the connectivity. This allows you to mitigate risk from exposure to non-conformant hardware and software in your fabrics.

You can click **+Add Ignore Host** to enter one or more destination IP addresses for which traffic will skip using the proxy.

The proxy server must permit these URLs:

```
svc.intersight.com
svc-static1.intersight.com
svc-static1.ucs-connect.com
```

If you do not want to configure a proxy, click **Skip Proxy** then click **Confirm**.

- h) (Optional) If your proxy server requires authentication, put a check in the **Authentication required for Proxy** checkbox and enter the login credentials.
- i) (Optional) Expand the **Advanced Settings** category and change the settings if required.

Under advanced settings, you can configure these settings:

- **App Network**—The address space used by the application's services running in the Nexus Dashboard. Enter the IP address and netmask.
- **Service Network**—An internal network used by Nexus Dashboard and its processes. Enter the IP address and netmask.
- **App Network IPv6**—If you put a check in the **Enable IPv6** checkbox earlier, enter the IPv6 subnet for the app network.
- **Service Network IPv6**—If you put a check in the **Enable IPv6** checkbox earlier, enter the IPv6 subnet for the service network.

For more information about the application and service networks, see [General prerequisites and guidelines, on page 11](#).

- j) Click **Next**.

Step 14

In the **Node Details** page, update the first node's information.

You have defined the Management network and IP address for the node into which you are currently logged in during the initial node configuration in earlier steps, but you must also enter the Data network information for the node before you can proceed with adding the other `primary` nodes and creating the cluster.

- a) For **Cluster Connectivity**, if your cluster is deployed in L3 mode, choose **BGP**. Otherwise, choose **L2**.

BGP configuration is required for the persistent IP addresses feature used by telemetry. This feature is described in more detail in the [BGP configuration and persistent IP addresses, on page 46](#) and [Nexus Dashboard persistent IP addresses, on page 40](#) sections.

Note

You can enable BGP at this time or in the Nexus Dashboard GUI after the cluster is deployed. All remaining nodes need to configure BGP if it is configured. You must enable BGP now if the data network of nodes have different subnets.

- b) Click the **Edit** button next to the first node.

The node's **Serial Number**, **Management Network** information, and **Type** are automatically populated, but you must enter the other information.

- c) For **Name**, enter a name for the node.

The node's **Name** will be set as its hostname, so it must follow the [RFC-1123](#) requirements.

Note

If you need to change the name but the **Name** field is not editable, run the CIMC validation again to fix this issue.

- d) For **Type**, choose **Primary**.

The first nodes of the cluster must be set to **Primary**. You will add the secondary nodes in a later step if required for higher scale.

- e) In the **Data Network** area, enter the node's data network information.

Enter the data network IP address, netmask, and gateway. Optionally, you can also enter the VLAN ID for the network. Leave the VLAN ID field blank if your configuration does not require VLAN. If you chose **BGP** for **Cluster Connectivity**, enter the ASN.

If you enabled IPv6 functionality in a previous page, you must also enter the IPv6 address, netmask, and gateway.

Note

If you want to enter IPv6 information, you must do so during the cluster bootstrap process. To change the IP address configuration later, you would need to redeploy the cluster.

All nodes in the cluster must be configured with either only IPv4, only IPv6, or dual stack IPv4/IPv6.

- f) If you chose **BGP** for **Cluster Connectivity**, then in the **BGP peer details** area, enter the peer's IPv4 address and ASN.

You can click + **Add IPv4 BGP peer** to add additional peers.

If you enabled IPv6 functionality in a previous page, you must also enter the peer's IPv6 address and ASN.

- g) Click **Save** to save the changes.

Step 15

In the **Node Details** screen, click **Add Node** to add the second node to the cluster.

If you are deploying a single-node cluster, skip this step.

Edit Node



General

Name *

Serial Number *

Type *

Management Network ⓘ

IPv4 Address/Mask *

IPv4 Gateway *

IPv6 Address/Mask

IPv6 Gateway

Data Network ⓘ

IPv4 Address/Mask *

IPv4 Gateway *

IPv6 Address/Mask

IPv6 Gateway

VLAN ⓘ

Enable BGP ☐

- a) In the **Deployment Details** area, provide the **Management IP Address** and **Password** for the second node

You defined the management network information and the password during the initial node configuration steps.

- b) Click **Validate** to verify connectivity to the node.

The node's **Serial Number** and the **Management Network** information are automatically populated after connectivity is validated.

- c) Provide the **Name** for the node.
d) From the **Type** dropdown, select `Primary`.

The first 3 nodes of the cluster must be set to `Primary`. You will add the secondary nodes in a later step if required for higher scale.

- e) In the **Data Network** area, provide the node's **Data Network** information.

You must provide the data network IP address, netmask, and gateway. Optionally, you can also provide the VLAN ID for the network. For most deployments, you can leave the VLAN ID field blank.

If you had enabled IPv6 functionality in a previous screen, you must also provide the IPv6 address, netmask, and gateway.

Note

If you want to provide IPv6 information, you must do it during cluster bootstrap process. To change IP configuration later, you would need to redeploy the cluster.

All nodes in the cluster must be configured with either only IPv4, only IPv6, or dual stack IPv4/IPv6.

- f) (Optional) If your cluster is deployed in L3 mode, **Enable BGP** for the data network.

BGP configuration is required for the persistent IP addresses feature. This feature is described in more detail in [BGP configuration and persistent IP addresses, on page 46](#) and the "Persistent IP Addresses" sections of the [Cisco Nexus Dashboard User Guide](#).

Note

You can enable BGP at this time or in the Nexus Dashboard GUI after the cluster is deployed.

If you choose to enable BGP, you must also provide the following information:

- **ASN** (BGP Autonomous System Number) of this node.
You can configure the same ASN for all nodes or a different ASN per node.
- For IPv6-only, the **Router ID** of this node.
The router ID must be an IPv4 address, for example `1.1.1.1`
- **BGP Peer Details**, which includes the peer's IPv4 or IPv6 address and peer's ASN.

- g) Click **Save** to save the changes.
h) Repeat this step for the final (third) primary node of the cluster.

Step 16

In the **Node Details** page, verify the information that you entered, then click **Next**.

Step 17

Choose the **Deployment Mode** for the cluster.

- a) Click **Add Persistent Service IPs/Pools** to provide the required persistent IP addresses.

For more information about persistent IP addresses, see the [Nexus Dashboard persistent IP addresses, on page 40](#) section.

- b) Click **Next** to proceed.

Step 18 In the **Summary** screen, review and verify the configuration information and click **Save** to build the cluster.

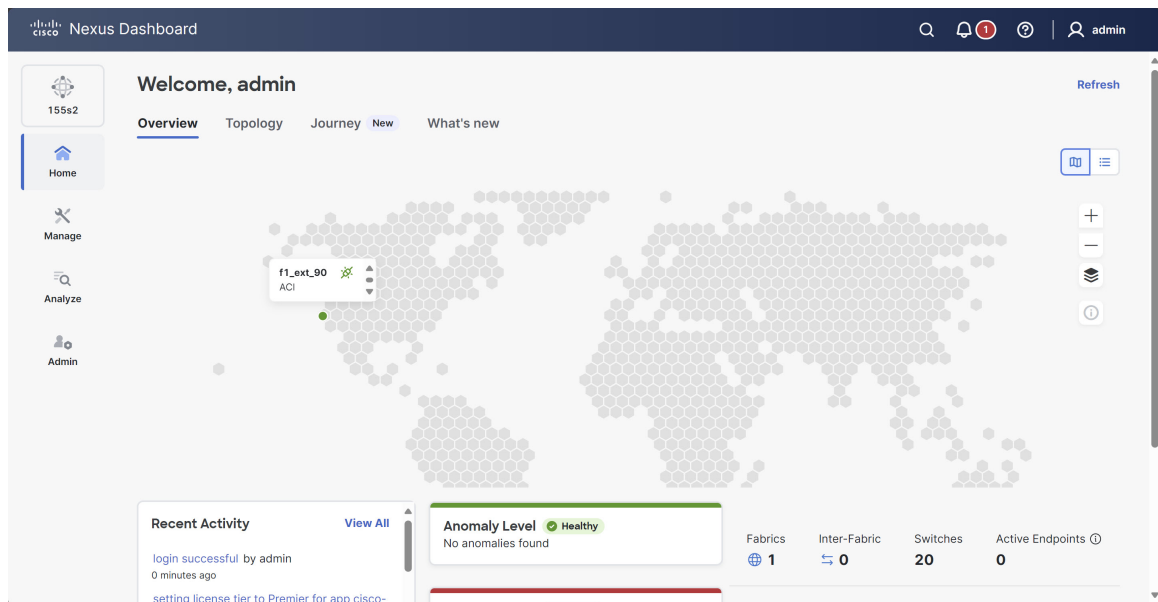
During the node bootstrap and cluster bring-up, the overall progress as well as each node's individual progress will be displayed in the UI. If you do not see the bootstrap progress advance, manually refresh the page in your browser to update the status.

It may take up to 30 minutes for the cluster to form and all the services to start. When cluster configuration is complete, the page will reload to the Nexus Dashboard GUI.

Step 19 Verify that the cluster is healthy.

After the cluster becomes available, you can access it by browsing to any one of your nodes' management IP addresses. The default password for the `admin` user is the same as the `rescue-user` password you chose for the first node. During this time, the UI will display a banner at the top stating "Service Installation is in progress, Nexus Dashboard configuration tasks are currently disabled".

After all the cluster is deployed and all services are started, you can look at the **Anomaly Level** on the **Home > Overview** page to ensure the cluster is healthy:



Alternatively, you can log in to any one node using SSH as the `rescue-user` using the password you entered during node deployment and using the `acs health` command to see the status:

- While the cluster is converging, you may see the following output:

```
$ acs health
k8s install is in-progress

$ acs health
k8s services not in desired state - [...]

$ acs health
k8s: Etcd cluster is not ready
```

- When the cluster is up and running, the following output will be displayed:

```
$ acs health
All components are healthy
```

Note

In some situations, you might power cycle a node (power it off and then back on) and find it stuck in this stage:

```
deploy base system services
```

This is due to an issue with `etcd` on the node after a reboot of the physical Nexus Dashboard cluster.

To resolve the issue, enter the `acs reboot clean` command on the affected node.

Step 20 (Optional) Connect your Cisco Nexus Dashboard cluster to Cisco Intersight for added visibility and benefits. Refer to [Working with Cisco Intersight](#) for detailed steps.

Step 21 After you have deployed Nexus Dashboard, see the [collections page](#) for this release for configuration information.

What to do next

The next task is to create the fabrics and fabric groups. See the *Creating Fabrics and Fabric Groups* article for this release on the [Cisco Nexus Dashboard collections page](#).



CHAPTER 8

Deploying a Virtual Nexus Dashboard (vND) in Nutanix

- [Prerequisites and guidelines for deploying vNDs in Nutanix, on page 169](#)
- [About the deployment script for QCOW2 images, on page 171](#)
- [Install the Nexus Dashboard cluster in Nutanix, on page 180](#)

Prerequisites and guidelines for deploying vNDs in Nutanix

You can now deploy virtual Nexus Dashboard (vNDs) on Nutanix Hyperconverged Infrastructure (HCI).

Before you proceed with deploying vNDs in Nutanix, you must follow these prerequisites and guidelines.

- vNDs on Nutanix supports both single-node and three-node deployments. The underlying Nutanix cluster can be of any size, provided that it has sufficient resources to host the required nodes and workload.
- You can enable the Controller, Telemetry, and Orchestration features. See the [Cisco Nexus Dashboard Verified Scalability Guide, Release 4.2.1](#) for more information.
- Scale support and co-hosting vary based on the cluster form factor. Use the [Nexus Dashboard Capacity Planning](#) tool to verify that the virtual form factor satisfies your deployment requirements.
- Review and complete the general prerequisites as described in [Prerequisites and Guidelines, on page 11](#).
- Management and data interfaces must be created with subnets of the same type, where both subnets are either `VLAN_Basic` or `VLAN`. Deployment of vND in Nutanix will fail if different subnet types are used for management and data interfaces.
- The CPU family used for the Nexus Dashboard VMs must support the AVX instruction set.
- The Nutanix node must have enough system resources, and the Nutanix system storage device must have I/O latency under 20ms.
 - See [Understanding system resources, on page 171](#) for more information on system resources. Only 1-node (data) and 3-node (data) virtual profiles are supported. You must explicitly reserve these resources on each node to achieve maximum performance and reliability.
 - See [Verify the I/O latency of a Nutanix storage device, on page 170](#) for more information on I/O latency.
- vND on Nutanix deployments are supported for LAN and SAN deployments.

- Nexus Dashboard should only be deployed on Nutanix Hyperconverged Infrastructure (HCI). For more information, see <https://www.cisco.com/site/us/en/products/computing/hyperconverged/nutanix/index.html>.
- Supported Nutanix software releases:

Acropolis operating system (AOS)	Acropolis Hypervisor (AHV)	Prism Central
6.10.x	20230302.x	2024.2.x
7.3.x	10.3.x	7.3.x

- We highly recommend that you deploy each Nexus Dashboard node on a separate Acropolis Hypervisor Host.
- There is no support for live-migration (this should be disabled on the Nexus Dashboard VMs).

Verify the I/O latency of a Nutanix storage device



Note Refer to [Performance benchmarking with Fio on Nutanix](#) for additional useful information.

When you deploy a Nexus Dashboard cluster on Nutanix Hyperconverged Infrastructure (HCI), the Nutanix storage device must have a latency under 20ms. This is done in the Performance Validation step if you use the `kvm_deployer` deployment script for QCOW2 images, as described in [About the deployment script for QCOW2 images, on page 149](#).

Follow these steps if you want to manually verify the I/O latency of a Nutanix storage device.

Procedure

-
- Step 1** Install the fio packet:
- ```
sudo apt-get install fio
```
- Step 2** Create a test-data folder.
- For example, create a folder named `test-data`:
- ```
# mkdir test-data
```
- Step 3** Run the Flexible I/O tester (FIO) test command.
- ```
fio --rw=write --ioengine=sync --fdatasync=1 --directory=test-data --size=22m --bs=2300 --name=mytest
```
- Step 4** After you use the command, confirm that the `99.00th=[value]` in the `fsync/fdatasync/sync_file_range` section is under 20ms.
-

## Understanding system resources

When deploying a Nexus Dashboard cluster in Nutanix, you must verify that you have enough system resources. There are multiple form factors supported when deploying a Nexus Dashboard cluster in Nutanix, and the amount of system resources needed for each node differs based on the form factor.



**Note** Deployment of virtual Nexus Dashboard (vNDs) in Nutanix is supported only on data nodes, as shown below. Deployment of vNDs in Nutanix is not supported on app nodes.

*Table 27: Per node resource requirements*

| Form factor       | Number of vCPUs | RAM    | Disk size | Cores per CPU |
|-------------------|-----------------|--------|-----------|---------------|
| 1-node vND (data) | 32              | 128 GB | 3077 GB   | 1             |
| 3-node vND (data) | 32              | 128 GB | 3077 GB   | 1             |

You will need to know the information above for your form factor when you go through the procedures in [Install the Nexus Dashboard cluster in Nutanix, on page 180](#).

## About the deployment script for QCOW2 images

Nexus Dashboard supports KVM virtualization across various environments, including RHEL and Nutanix. Nexus Dashboard deployments are typically done using QCOW2 images; however, the diversity of KVM variants and host configurations creates challenges in tracking supported versions and ensuring optimal device emulation for performance.

The deployment script for QCOW2 images addresses these challenges by providing a deployment validation and orchestration script that is packaged with the QCOW2 image. This script is designed to streamline the deployment process, ensure compatibility, and help maintain consistent performance standards across supported platforms.

### Key Features

These are the key features for the deployment script for QCOW2 images:

- **Platform Compatibility Check:** Verifies that the underlying host operating system is a supported distribution and version.
- **Resource Validation:** Assesses whether the host meets minimum compute requirements, including CPU, memory, and storage resources.
- **Performance Validation:**
  - **Disk I/O and Network Latency:** Measures key performance indicators such as disk I/O and network I/O latency, ensuring they fall within qualified ranges.
  - **Device Emulation:** Confirms that the correct device emulation (such as network interfaces and virtual disk devices) is available for optimal performance.

- **Deployment Orchestration:** Automates the deployment process for both single-node and multi-node clusters, orchestrating virtual machines with the appropriate device mappings (such as network and storage devices).

Because storage is the most common bottleneck, the script categorizes storage into "Performance Tiers":

- **Disk Type & Placement (SSD/NVMe):** Validates for high-IOPs workloads.
- **RAID vs. JBOD:** Validates write speed.
- **Remote Storage (iSCSI/FC/NFS):** Performs a "Pre-Flight Latency Check":
  - **Sync Latency:** Must be < 10ms for standard apps and < 1ms for high-performance databases.
  - **IOPS Validation:** Uses a temporary `fio` test to ensure the required IOS.

## Prerequisites for the deployment script for QCOW2 images

These are the prerequisites before using the deployment script for QCOW2 images.

| Requirement           | Details                                                                             |
|-----------------------|-------------------------------------------------------------------------------------|
| Nutanix Prism Central | Reachable on port 9440 (HTTPS) from the machine running the deployer                |
| Nutanix Prism Element | Reachable on port 9440 (HTTPS); used for storage, networking, and VM console access |
| Prism credentials     | A user account with permissions to create images, list subnets, and manage VMs      |
| ND qcow2 image        | A local path to the Nexus Dashboard .qcow2 image file                               |
| Network info          | Subnet name(s) configured in Prism, and 3 management IPs + a gateway                |
| Go toolchain          | Go 1.22+ to build from source (or use a pre-built binary)                           |

### Understanding the differences between Prism Central and Prism Element

The deployment script uses **Prism Central** (PC) for all API operations — image uploads, VM creation, and lifecycle management. You will need its IP and credentials for the configuration file.

**Prism Element** (PE) manages the individual Nutanix cluster and is useful for:

- Viewing storage containers and disk utilization
- Checking network/subnet configuration at the cluster level
- Accessing the VM serial console for troubleshooting boot issues
- Monitoring host-level health and hardware status

| Interface     | URL                                        | When you need it                                                                  |
|---------------|--------------------------------------------|-----------------------------------------------------------------------------------|
| Prism Central | <code>https://prism_central_ip:9440</code> | Deploying, managing, and destroying VMs; uploading images; verifying subnet names |
| Prism Element | <code>https://prism_element_ip:9440</code> | Checking storage containers, network config, VM console access, host health       |

## Use the deployment script for QCOW2 images

### Procedure

**Step 1** Locate the tar file with the deployment script.

The tar file should be available at the same location as the QCOW2 image. For example, in the [Software Download site for Nexus Dashboard](#) for the 4.3.1 release and later, you should see the tar file with the deployment script for QCOW2 images at the same area as the QCOW2 image itself, such as this:

```
nd-dk9.4.3.1.175.qcow2.tar.gz
```

**Step 2** Download the tar file with the deployment script.

**Step 3** Extract the tar file.

For example:

```
tar -xzf nd-dk9.4.3.1.175.qcow2.tar.gz
```

You should see output similar to the following after extracting the tar file and listing the contents in the directory:

```
#ls -lrt
total 24730576
-rwxr-xr-x 1 root root 35991736 May 5 23:35 kvm_deployer
-rw-r--r-- 1 root root 12651498496 May 17 03:44 nd-dk9.4.3.1.175.qcow2
-rw-r--r-- 1 root root 12636610743 May 18 22:44 nd-dk9.4.3.1.175.qcow2.tar.gz
```

**Step 4** Use the **kvm\_deployer config** command to create a yaml file to use with the deployment script.

Make the appropriate choices at each stage based on your deployment type.

- For example, to generate a one-node yaml file:

```
./kvm_deployer config
Configure Nexus Dashboard deployment:
Deployment Platform (kvm/nutanix): nutanix
Image Path (qcow2): /home/nd-base/nd-dk9.4.3.1.175.qcow2
Cluster Size: 1
VM Type (app/data/large/xlarge): data
Admin Password:
Configure Nutanix platform settings:
Prism Central IP: 192.168.94.164
Nutanix Username: admin
Nutanix Password:
Cluster Name: sit_reg1_nutanix
Network Name: mgmt
VCPUs per Socket (default: 2): (current: 0): 2
```

```

Configure VM (1) for Nutanix:
 VM Name (lowercase, numbers, dashes, max 63 chars): node1
 Management Network (e.g., 192.168.1.100/24): 192.168.94.166/24
 Management Gateway IP: 192.168.94.1
[] Configuration Summary:
Deployment Platform : nutanix
Image Path (qcow2) : /home/nd-base/nd-dk9.4.3.1.175.qcow2
Admin Password : [hidden]
VM Type : data
Nutanix Configuration:
 Prism Central IP : 192.168.94.164
 Username : admin
 Password : [hidden]
 Cluster Name : sit_reg1_nutanix
 Network Name : mgmt
 VCPUs per Socket : 2
VMs:
 VM Name : node1
 Management Network : 192.168.94.166/24
 Management Gateway IP : 192.168.94.1
Proceed with current config? (y/n): y
File name to save configuration: nutanix_1node_deploy.yaml

```

- To generate a three-node yaml file:

```

./kvm_deployer config
Configure Nexus Dashboard deployment:
 Deployment Platform (kvm/nutanix): nutanix
 Image Path (qcow2): /home/nd-base/nd-dk9.4.3.1.175.qcow2
 Cluster Size: 3
 VM Type (app/data/large/xlarge): data
 Admin Password:
Configure Nutanix platform settings:
 Prism Central IP: 192.168.94.164
 Nutanix Username: admin
 Nutanix Password:
 Cluster Name: sit_reg1_nutanix
 Network Name: mgmt
 VCPUs per Socket (default: 2): (current: 0): 2
Configure VM (1) for Nutanix:
 VM Name (lowercase, numbers, dashes, max 63 chars): node1
 Management Network (e.g., 192.168.1.100/24): 192.168.94.166/24
 Management Gateway IP: 192.168.94.1
Configure VM (2) for Nutanix:
 VM Name (lowercase, numbers, dashes, max 63 chars): node2
 Management Network (e.g., 192.168.1.100/24): 192.168.94.167/24
 Management Gateway IP: 192.168.94.1
Configure VM (3) for Nutanix:
 VM Name (lowercase, numbers, dashes, max 63 chars): node3
 Management Network (e.g., 192.168.1.100/24): 192.168.94.168/24
 Management Gateway IP: 192.168.94.1
[] Configuration Summary:
Deployment Platform : nutanix
Image Path (qcow2) : /home/nd-base/nd-dk9.4.3.1.175.qcow2
Admin Password : [hidden]
VM Type : data
Nutanix Configuration:
 Prism Central IP : 192.168.94.164
 Username : admin
 Password : [hidden]
 Cluster Name : sit_reg1_nutanix
 Network Name : mgmt
 VCPUs per Socket : 2
VMs:

```

```

VM Name : node1
Management Network : 192.168.94.166/24
Management Gateway IP : 192.168.94.1
VM Name : node2
Management Network : 192.168.94.167/24
Management Gateway IP : 192.168.94.1
VM Name : node3
Management Network : 192.168.94.168/24
Management Gateway IP : 192.168.94.1
Proceed with current config? (y/n): y
File name to save configuration: nutanix_3node_deploy.yaml

```

**Step 5** (Optional) Use a text editor to verify the information saved in the yaml file, if necessary.

For example, for the three-node example above, the yaml file would have information similar to the following:

```

cat nutanix_3node_deploy.yaml
deployment_type: nutanix
image_path: /var/www/html/ND_IMAGES/nd-dk9.4.3.1.75.qcow2
cluster_size: "3"
vms:
- name: node1
 disk_path: ""
 management_network: 192.168.94.166/24
 management_gateway: 192.168.94.1
 management_bridge: ""
 data_bridge: ""
 host_ip: ""
 host_user: ""
 host_password: ""
- name: node2
 disk_path: ""
 management_network: 192.168.94.167/24
 management_gateway: 192.168.94.1
 management_bridge: ""
 data_bridge: ""
 host_ip: ""
 host_user: ""
 host_password: ""
- name: node3
 disk_path: ""
 management_network: 192.168.94.168/24
 management_gateway: 192.168.94.1
 management_bridge: ""
 data_bridge: ""
 host_ip: ""
 host_user: ""
 host_password: ""
flavor: data
admin_password:
nutanix:
 prism_central_ip: 192.168.94.164
 username: admin
 password:
 cluster_name: sit_reg1_nutanix
 vcpus_per_socket: 2
 network_name: mgmt
#

```

**Step 6** Run the deployment script, using the yaml file as the input.

```
./kvm_deployer deploy --config nutanix_3node_deploy.yaml
```

Output similar to the following displays:

```

Deploying VM node1 using kvm
Cleanup old images
[] Cleaned image /home/nd-node1/boot.img
[] Cleaned image /home/nd-node1/disk.img
[] Cleaned image /home/nd-node1/config-drive.img
Creating boot disk, it can take several minutes, pls wait
[] Boot disk created successfully
Creating raw disk, it can take several minutes, pls wait
[] Raw disk created successfully
Creating configuration drive
[] Configuration drive created successfully

Starting VM
[] VM started successfully
[] KVM VM: node1 deployed successfully
Deploying VM node2 using kvm
Copying file /home/nd-base/nd-dk9.4.3.1.175.qcow2 to /home/nd-node2/nd-dk9.4.3.1.175.qcow2
[] File /home/nd-base/nd-dk9.4.3.1.175.qcow2 copied to /home/nd-node2/nd-dk9.4.3.1.175.qcow2
Cleanup old images
[] Cleaned image /home/nd-node2/boot.img
[] Cleaned image /home/nd-node2/disk.img
[] Cleaned image /home/nd-node2/config-drive.img
Creating boot disk, it can take several minutes, pls wait
[] Boot disk created successfully
Creating raw disk, it can take several minutes, pls wait
[] Raw disk created successfully
Creating configuration drive
Copying file /home/nd-node1/node2-bootstrap.json to /home/nd-node2/apic-sn-cfg-mount/bootstrap.json

[] File /home/nd-node1/node2-bootstrap.json copied to /home/nd-node2/apic-sn-cfg-mount/bootstrap.json

[] Configuration drive created successfully
Starting VM
[] VM started successfully
[] KVM VM: node2 deployed successfully
Deploying VM node3 using kvm
Copying file /home/nd-base/nd-dk9.4.3.1.175.qcow2 to /home/nd-node3/nd-dk9.4.3.1.175.qcow2
[] File /home/nd-base/nd-dk9.4.3.1.175.qcow2 copied to /home/nd-node3/nd-dk9.4.3.1.175.qcow2
Cleanup old images
[] Cleaned image /home/nd-node3/boot.img
[] Cleaned image /home/nd-node3/disk.img
[] Cleaned image /home/nd-node3/config-drive.img
Creating boot disk, it can take several minutes, pls wait
[] Boot disk created successfully
Creating raw disk, it can take several minutes, pls wait
[] Raw disk created successfully
Creating configuration drive
Copying file /home/nd-node1/node3-bootstrap.json to /home/nd-node3/apic-sn-cfg-mount/bootstrap.json

[] File /home/nd-node1/node3-bootstrap.json copied to /home/nd-node3/apic-sn-cfg-mount/bootstrap.json

[] Configuration drive created successfully
Starting VM
[] VM started successfully
[] KVM VM: node3 deployed successfully

```

When you run deploy, the tool performs the following steps for each of the VMs:

- a. Image upload — Checks if the qcow2 image already exists in Prism Central (matched by filename). If not, uploads it and waits until it becomes active.
- b. Subnet resolution — Resolves network\_name (and data\_network\_name if set) to Prism subnet UUIDs.
- c. VM creation — Creates the VM in a powered-off state with:

- A boot disk cloned from the uploaded image
- A secondary 500 GiB raw data disk
- NIC(s) attached to the resolved subnet(s)
- Cloud-init user data carrying the admin password, management IP, and gateway

d. Power on — Powers on the VM and waits for the task to complete.

**Step 7** Complete the post-deployment steps.

- a) **Cluster formation** — Once all the nodes are running, complete the Nexus Dashboard cluster setup through the Nexus Dashboard UI using the management IP addresses from your config.

The screenshot shows the Cisco Nexus Dashboard interface for the 'Cluster Bringup' configuration. The browser address bar shows 'https://1/cluster-bringup'. The page has a dark blue header with the Cisco logo and 'Nexus Dashboard' text, and a user profile 'admin' in the top right. A left sidebar contains a 'Journey' section with icons for 'Admin' and 'Cluster Bringup'. The main content area is titled 'Cluster Bringup' and features a vertical list of steps: 1. Basic information (selected), 2. Configuration, 3. Node details, 4. Persistent IPs, and 5. Summary. The 'Basic information' step is expanded, showing a warning message: 'Cluster name cannot be changed after Cluster Bringup. Changing this will require a cluster rebuild.' Below this, there is a text input field for 'Cluster name' containing 'Nd-KVM-133'. Underneath, the 'Select the Nexus Dashboard implementation type' section has two radio buttons: 'LAN' (selected) and 'SAN'. The 'LAN' option includes the text 'Includes Cisco NX-OS, ACI, IOS XE, IOS XR and other 3rd party use cases', while the 'SAN' option includes 'Includes fiber channel and storage use cases'. At the bottom of the configuration panel are 'Cancel' and 'Next' buttons.

## Use the deployment script for QCOW2 images

## Cluster connectivity \* ⓘ

- ☒ L2
- ☐ BGP

| Serial number | Name      | Type    | Management network                   | Data network                                     | Configuration status ⓘ |                                                                                                                                                                         |
|---------------|-----------|---------|--------------------------------------|--------------------------------------------------|------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 873DF24266F2  | node1-134 | Primary | IPv4 address:<br>IPv4 gateway:<br>   | IPv4 address:<br>/24<br>IPv4 gateway:<br>VLAN: - | Successful             |   |
| 4EE78A04C3CE  | node2-135 | Primary | IPv4 address:<br>IPv4 gateway:<br>   | IPv4 address:<br>/24<br>IPv4 gateway:<br>VLAN: - | Successful             |   |
| WVT2922012Y   | node3-136 | Primary | IPv4 address:<br>IPv4 gateway:<br>.1 | IPv4 address:<br>/24<br>IPv4 gateway:<br>VLAN: - | Successful             |   |

+ Add node

Cancel

Back

Validate

← → ↻ Not Secure https://cluster-bringup

Nexus Dashboard

Cluster Bringup

Summary

Note: Once installation is complete, Nexus Dashboard will reboot and you will need to log back in.

The following items cannot be changed after Cluster Bringup: Cluster name, App subnet, and Node name. Changing these will require a cluster rebuild.

Basic information

Cluster name: NSD-KVSD-133

Persona: L2N

DNS configuration

# DNS provider IP address

1

# DNS search domain

1: cisco.com

NTP setup

NTP authentication: Enabled

# NTP hostname or IP address

Key: Preferred

1: Yes

Proxy

Proxy configuration: Skip

Advanced settings

App network

Service network

Node details

Cluster mode: L2

| Serial number | Name    | Type                               | Management network                        | Data network |
|---------------|---------|------------------------------------|-------------------------------------------|--------------|
| node2-135     | Primary | IPv4 address:<br>IPv4 gateway:<br> | IPv4 address:<br>IPv4 gateway:<br>VLAN: - |              |
| node1-134     | Primary | IPv4 address:<br>IPv4 gateway:<br> | IPv4 address:<br>IPv4 gateway:<br>VLAN: - |              |
| node3-136     | Primary | IPv4 address:<br>IPv4 gateway:<br> | IPv4 address:<br>IPv4 gateway:<br>VLAN: - |              |

Cancel

Back Save

## Optional post-deployment tasks

### Troubleshooting

| Symptom                      | Check                                                                                                                                                             |
|------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Connection refused / timeout | Verify Prism Central is reachable on port 9440 from your machine                                                                                                  |
| Authentication failure       | Verify <code>username</code> and <code>password</code> in the config                                                                                              |
| Image upload hangs           | Ensure <code>image_path</code> points to a valid qcow2 file and there is sufficient storage in Prism                                                              |
| Subnet not found             | Verify <code>network_name</code> matches the exact subnet name shown in Prism Central                                                                             |
| VM not getting IP            | Check that cloud-init <code>management_network</code> uses valid CIDR notation and the gateway is correct                                                         |
| VM stuck at boot             | Use Prism Element ( <a href="https://&lt;prism_element_ip&gt;:9440">https://&lt;prism_element_ip&gt;:9440</a> ) to open the VM serial console and check boot logs |
| Storage issues               | Check available storage containers and disk space in Prism Element                                                                                                |

### Destroy the deployment

If you want to destroy the deployment, enter this command:

```
./kvm_deployer destroy --config <yaml-filename>.yaml
```

For example:

```
$./kvm_deployer destroy --config nutanix_3node_deploy.yaml
Destroying VM node1 using nutanix
Destroying VM node1 using nutanix
Looking for VM: node1
Found VM node1 with UUID: 94fe9518-5196-4807-a6cc-183dlabe760e
Powering off VM: 94fe9518-5196-4807-a6cc-183dlabe760e
VM power off request submitted
Deleting VM: 94fe9518-5196-4807-a6cc-183dlabe760e
VM deletion task submitted: 8178938b-5014-474b-b168-34c3bf5b1577
Waiting for task completion: 8178938b-5014-474b-b168-34c3bf5b1577
Task status: RUNNING, progress: 0%
Task status: SUCCEEDED, progress: 100%
[] Task completed successfully (status: SUCCEEDED)
[] VM deleted successfully
[] VM node1 destroyed successfully
Destroying VM node2 using nutanix
Destroying VM node2 using nutanix
Looking for VM: node2
Found VM node2 with UUID: f1424251-182d-4518-9acc-f06ba89d4883
Powering off VM: f1424251-182d-4518-9acc-f06ba89d4883
VM power off request submitted
Deleting VM: f1424251-182d-4518-9acc-f06ba89d4883
VM deletion task submitted: a8d41687-3d84-4a38-91cc-5dfce0688351
Waiting for task completion: a8d41687-3d84-4a38-91cc-5dfce0688351
Task status: RUNNING, progress: 0%
Task status: SUCCEEDED, progress: 100%
[] Task completed successfully (status: SUCCEEDED)
```

```

[] VM deleted successfully
[] VM node2 destroyed successfully
Destroying VM node3 using nutanix
Destroying VM node3 using nutanix
Looking for VM: node3
Found VM node3 with UUID: 8cdc2037-fc83-43af-90f6-869c00aa2fdc
Powering off VM: 8cdc2037-fc83-43af-90f6-869c00aa2fdc
VM power off request submitted
Deleting VM: 8cdc2037-fc83-43af-90f6-869c00aa2fdc
VM deletion task submitted: 710c58a7-0728-4c4e-b63a-f6d90e3b2669
Waiting for task completion: 710c58a7-0728-4c4e-b63a-f6d90e3b2669
Task status: RUNNING, progress: 0%
Task status: SUCCEEDED, progress: 100%
[] Task completed successfully (status: SUCCEEDED)
[] VM deleted successfully
[] VM node3 destroyed successfully
$

```

# Install the Nexus Dashboard cluster in Nutanix

## Procedure

### Step 1

Download the Cisco Nexus Dashboard image.

- Navigate to the [Software Download](#) page.
- In the left navigation pane, select the Nexus Dashboard release you want to download.

Nutanix deployments are supported starting with Nexus Dashboard Release 4.2.

- Download the Cisco Nexus Dashboard image for Linux KVM (`nd-dk9.<version>.qcow2`).

Nutanix uses the same `qcow2` image as Linux KVM.

### Step 2

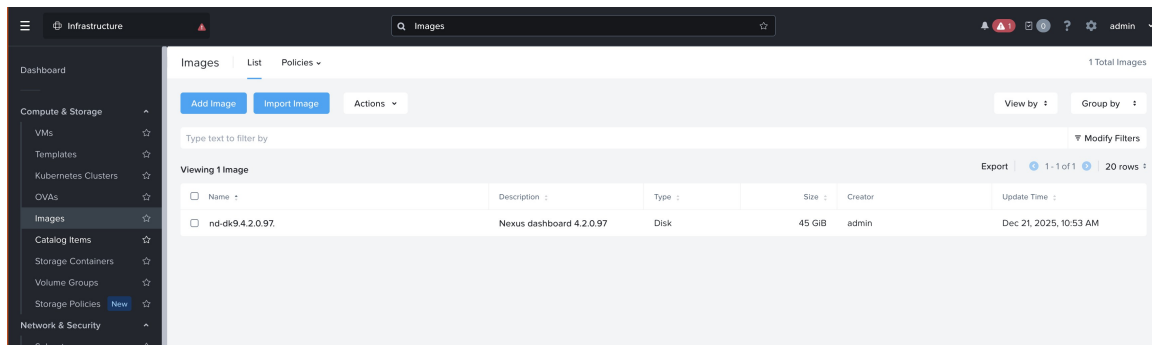
Copy the downloaded `qcow2` image to a web server accessible by Prism Central over HTTP or HTTPS.

### Step 3

Access the Prism Central UI.

### Step 4

Navigate to **Infrastructure > Compute & Storage > Images**, then click **Add Image**.



### Step 5

Click the **URL** option, then enter the URL path to the Nexus Dashboard `qcow2` image and click **Add URL**.

#### Note

Because the Nexus Dashboard `qcow2` image is larger than 2 GB, the **Image File** option cannot be used to add this image.

1

Select Image

2

Select Location

---

Image Source

☐ Image File ☒ URL ☐ VM Disk

Image URL

g://172.24.80.172:8082/artifactory/nd-unified/nd/release/v4.2.0.99/nd-dk9.4.2.0.99.qcow2

Authentication (optional) 

+ Add URL

Cancel

Next

**Step 6** When the image information appears in the window, click **Next**.

## Image Source

☐ Image File
 ☒ URL
 ☐ VM Disk

## Image URL

Authentication (optional) 

+ Add URL

Source: nd-dk9.4.2.0.99.qcow2

Remove Image

## General

Name

Type

Disk

Description

Checksum

SHA-1

Authentication (optional) 

Cancel

Next

## Step 7

Make these configurations in the **2. Select Location** window.

- In the **Placement Method** area, click **Place image directly on clusters**.
- Choose the clusters to use for placement, then click **Save**.

The image upload could take a few minutes, depending on the speed. Wait until the process is completed before proceeding to the next step.

✓ Select Image      2 Select Location

### Placement Method

☒ Place image directly on clusters

This option is good for smaller environments. The image will be placed on all selected clusters below.

☐ Place image using Image Placement policies

This option is good for larger environments. It requires you to first set up Image Placement policies between categories assigned to clusters and categories assigned to images. From there on, you only need to associate a relevant category to an image while uploading it here.

### Select Clusters

Select the set of clusters to use for placement

| <input checked="" type="checkbox"/> Name ↑      | Bandwidth Limit |
|-------------------------------------------------|-----------------|
| <input checked="" type="checkbox"/> vnd-nutanix |                 |

Back Cancel Save

## Step 8

Deploy the vND VM.

- Navigate To **Infrastructure > Compute & Storage > VMs**, then click **Create VM**.
- In the **1. Configuration** window, enter the CPU and Memory values required for Nexus Dashboard, then click **Next**.

Use the information provided in the [Understanding system resources, on page 171](#) for these values.

## Create VM

1

Configuration

2

Resources

3

Management

4

Review

Name

nd-1

Description

(Optional)

Cluster

sauroabh-nutanix

Number of VMs

1

VM Properties

CPU

32

vCPU

Cores Per CPU

1

Cores

Memory

128

GiB

Advanced Settings

Cancel

Next

- c) In the **2. Resources** window, click **Attach Disk** and attach the boot disk to the VM, then click **Save**.

## Create VM

✓ Configuration 2 Resources 3 Management 4 Review

**i** Any storage policy applied later, will manage the storage properties for all VM disks. Data placement will remain unaffected. [Learn More](#)



## Disks

Attach Disk

## Networks

Attach to Subnet

Want to use this VM as a Traffic Mirror Destination? [Add Mirror Destination NIC](#)

## Boot Configuration

☒ UEFI BIOS Mode

UEFI BIOS Mode supports enhanced Shield VM security settings.

☐ Legacy BIOS Mode

## Shield VM Security Settings

☐ Secure Boot

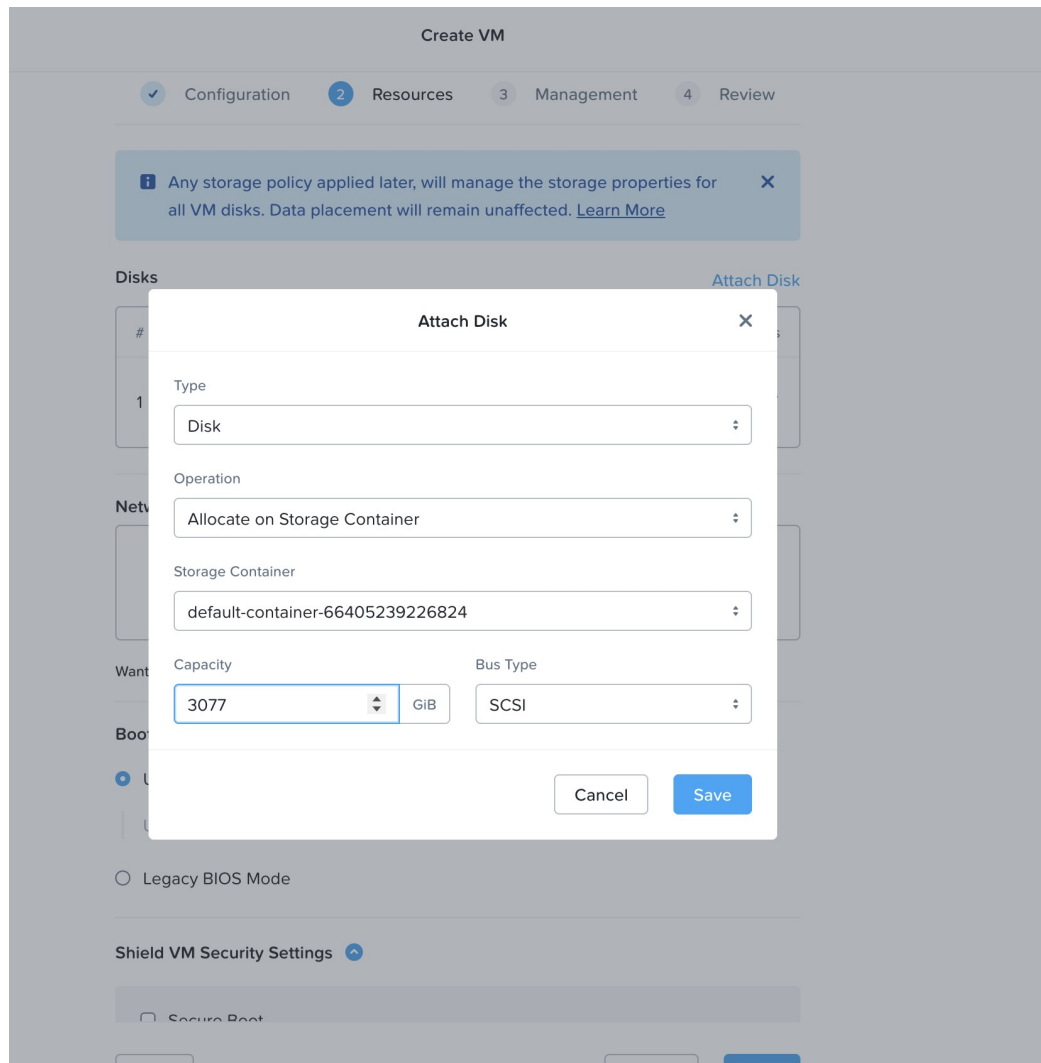
Back

Cancel

Next

The screenshot shows the 'Create VM' wizard in Nutanix, specifically the 'Resources' step. The wizard has four steps: Configuration, Resources (current), Management, and Review. A notification banner at the top states: 'Any storage policy applied later, will manage the storage properties for all VM disks. Data placement will remain unaffected. [Learn More](#)'. The 'Disks' section is visible, and the 'Attach Disk' dialog is open. The dialog has the following fields: 'Type' (set to 'Disk'), 'Operation' (set to 'Clone from Image'), 'Image' (set to 'nd-dk9.4.2.0.99.qcow2'), 'Capacity' (set to '45'), 'Bus Type' (set to 'SCSI'), and 'Unit' (set to 'GiB'). There are 'Cancel' and 'Save' buttons at the bottom right of the dialog. The background shows the 'Disks' section of the 'Create VM' form, with a 'Net' section partially visible.

- d) In the **2. Resources** window, click **Attach Disk** again and add a secondary disk.  
See [Understanding system resources, on page 171](#) for more information.



- e) In the **2. Resources** window, click **Attach to Subnet** and make the necessary configurations for the first subnet.

**Note**

Management and data interfaces must be created with subnets of the same type, where both subnets are either `VLAN Basic` or `VLAN`. Deployment of vND in Nutanix will fail if different subnet types are used for management and data interfaces.

## Create VM

☒ Configuration
 ☒ **Resources**
☐ 3 Management
 ☐ 4 Review

 Any storage policy applied later, will manage the storage properties for all VM disks. Data placement will remain unaffected. [Learn More](#)


## Disks

[Attach Disk](#)

| # | Type | Source                           | Size     | Bus Type | Actions                                                                                                                                                                 |
|---|------|----------------------------------|----------|----------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 1 | Disk | nd-dk9.4.2.0.99.qcow2<br>Image   | 45 GiB   | SCSI     |   |
| 2 | Disk | default-container-66405239226824 | 3077 GiB | SCSI     |   |

## Networks

Attach to Subnet

Want to use this VM as a Traffic Mirror Destination? [Add Mirror Destination NIC](#)

## Boot Configuration

☒ UEFI BIOS Mode

UEFI BIOS Mode supports enhanced Shield VM security settings.

☐ Legacy BIOS Mode

Back

Cancel

Next

- In the **Subnet** field, choose the management subnet.

- Configure the entry in the **Attachment type** field to reflect your environment.

Click **Save** when you have completed the configurations in this window.

Attach to Subnet

Subnet Attachment

Subnet

mgmt

|         |             |                |
|---------|-------------|----------------|
| VLAN ID | IPAM        | Virtual Switch |
| 0       | Not Managed | br0            |

Network Connection State

Connected

NIC Configuration

Attachment Type

Access

Cancel

Save

- f) In the **2. Resources** window, click **Attach to Subnet** and make the necessary configurations for the second subnet.
- In the **Subnet** field, choose the data subnet.
  - Configure the entry in the **Attachment type** field to reflect your environment.

Click **Save** when you have completed the configurations in this window.

**Create VM**

☒ Configuration   
 ☒ **Resources**   
 ☐ Management   
 ☐ Review

Any storage policy applied later, will manage the storage properties for all VM disks. Data placement will remain unaffected. Learn More

**Attach to Subnet** ✕

**Subnet Attachment**

Subnet

data-network

| VLAN ID | IPAM        | Virtual Switch |
|---------|-------------|----------------|
| 102     | Not Managed | br1            |

Network Connection State

Connected

**NIC Configuration** ^

Attachment Type

Access

Cancel
Save

Legacy BIOS Mode

Back
Cancel
Next

- g) In the **2. Resources** window, in the **Boot Configuration** area, choose **Legacy BIOS Mode** and confirm the choice.

## Create VM

all VM disks. Data placement will remain unaffected. [Learn More](#)

## Disks

[Attach Disk](#)

| # | Type | Source                           | Size     | Bus Type | Actions |
|---|------|----------------------------------|----------|----------|---------|
| 1 | Disk | nd-dk9.4.2.0.99.qcow2<br>Image   | 45 GiB   | SCSI     |         |
| 2 | Disk | default-container-66405239226824 | 3077 GiB | SCSI     |         |

## Networks

[Attach to Subnet](#)

| Subnet       | VLAN ID / VPC | Private IP | Public IP | Actions |
|--------------|---------------|------------|-----------|---------|
| mgmt         | 0             | None       | None      |         |
| data-network | 102           | None       | None      |         |

Want to use this VM as a Traffic Mirror Destination? [Add Mirror Destination NIC](#)

## Boot Configuration

☐ UEFI BIOS Mode

UEFI BIOS Mode supports enhanced Shield VM security settings.

☒ Legacy BIOS Mode

Set Boot Priority

Default Boot Order (CD-ROM, Disk, Network)

Back

Cancel

Next

h) Navigate through the **Management** window and click **Next** without making changes.

## Create VM

✓ Configuration
 ✓ Resources
 3 Management
 4 Review

Enable 'Default-Storage' policy to manage the storage configurations across all VM disks. The policy applies via category 'Storage:\$Default'.

☐ Enable 'Default-Storage' policy
 [How does it work?](#)

**i** Applies to VMs on clusters with AOS 6.1 or above only.

Categories

Type to search...

**i** Tag the VM with Category: Value to assign policies associated with value

Timezone

(UTC) UTC

**i** Use UTC timezone for Linux VMs and local timezone for Windows VMs.

☐ Use this VM as an Agent VM **?**

**Guest Customization**

Script Type

No Customization

Configuration Method

Custom Script

Back

Cancel

Next

i) Click **Create VM**.

Do not power on the VM after the deployment.

## Create VM

Instance Properties 32 vCPU, 1 Core, 128 GB

Memory Overcommit Disabled

Advanced processor compatibility -

Resources [Edit](#)

## Disks

| # | Type | Source                                                                                                  | Size     | Bus Type |
|---|------|---------------------------------------------------------------------------------------------------------|----------|----------|
| 1 | Disk | nd-dk9.4.2.0.99.qcow2  | 45 GiB   | SCSI     |
| 2 | Disk | default-container-66405239226824                                                                        | 3077 GiB | SCSI     |

## Networks

| Subnet       | VLAN ID / VPC | Private IP | Public IP |
|--------------|---------------|------------|-----------|
| mgmt         | 0             | None       | None      |
| data-network | 102           | None       | None      |

## Security

Boot Configuration Legacy BIOS Mode: Default Boot Order

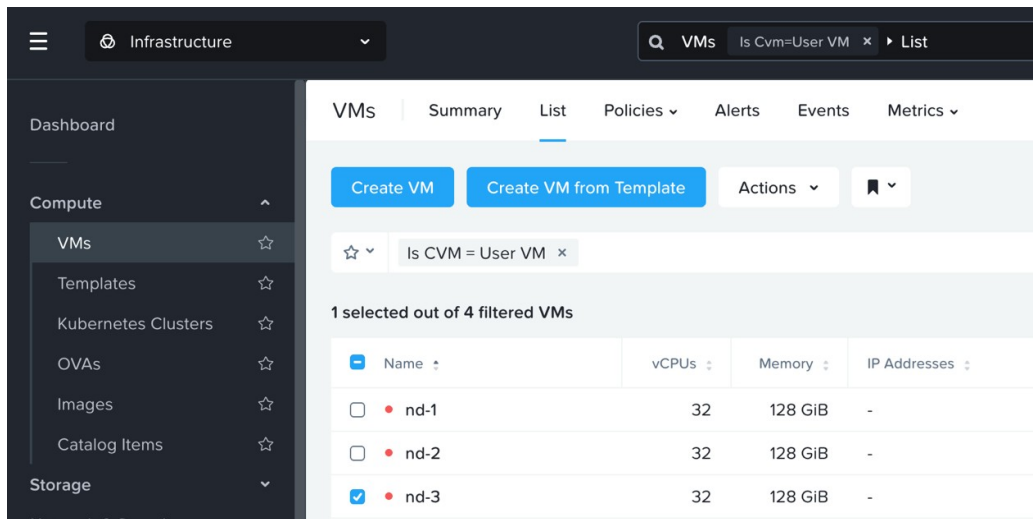
Management [Edit](#)

Categories None

Timezone UTC

[Back](#)[Cancel](#)[Create VM](#)

- j) If you are deploying in a 3-node cluster, repeat these steps for each Nexus Dashboard node in the cluster.



- k) Attach a serial console to a VM.

Nexus Dashboard requires a serial port for the configuration and bootstrap. To manage and add a new serial port, log on to any CVM CLI using SSH.

#### Note

See [Serial Console Redirection to a Telnet Port](#) for additional useful information.

```
nutanix@NTNX-35b35878-A-CVM:<nodeIP>:~$ accli vm.list
VM name VM UUID
PC-NameOption-1 201b8cf4-7ee6-409b-bb46-fca4265b6f70
nd-1 091c2d79-da94-4370-a77f-eccd998e2095
nd-2 58a04d9e-85da-4476-8735-bc6ed71e31ca
nd-3 c7f8feeb-b1c0-44a0-a5df-0dcfc30b07d8
nutanix@NTNX-35b35878-A-CVM:<nodeIP>:~$ accli vm.serial_port_create nd-1 index=0 type=kServer
VmUpdate: pending
VmUpdate: complete
lnutanix@NTNX-35b35878-A-CVM:<nodeIP>:~$ accli vm.serial_port_create nd-2 index=0 type=kServer
VmUpdate: pending
VmUpdate: complete
nutanix@NTNX-35b35878-A-CVM:<nodeIP>:~$ accli vm.serial_port_create nd-3 index=0 type=kServer
VmUpdate: pending
VmUpdate: complete
nutanix@NTNX-35b35878-A-CVM:<nodeIP>:~$
```

- l) Power on the Nexus Dashboard VMs.

## Step 9

Open one of the node's console and configure the node's basic information.

- a) Press any key to begin initial setup.

You will be prompted to run the first-time setup utility:

```
[OK] Started atomix-boot-setup.
Starting Initial cloud-init job (pre-networking)...
Starting logrotate...
Starting logwatch...
Starting keyhole...
[OK] Started keyhole.
[OK] Started logrotate.
[OK] Started logwatch.
```

Press any key to run first-boot setup on this console...

- b) Enter and confirm the `admin` password

This password will be used for the `rescue-user` SSH login as well as the initial GUI password.

**Note**

You must provide the same password for all nodes or the cluster creation will fail.

```
Admin Password:
Reenter Admin Password:
```

- c) Enter the management network information.

```
Management Network:
IP Address/Mask: <nodeIP>/<subnet>
Gateway: <gatewayIP>
```

- d) For the first node only, designate it as the "Cluster Leader".

You will log into the cluster leader node to finish configuration and complete cluster creation.

```
Is this the cluster leader?: y
```

- e) Review and confirm the entered information.

You will be asked if you want to change the entered information. If all the fields are correct, choose `n` to proceed. If you want to change any of the entered information, enter `y` to re-start the basic configuration script.

```
Please review the config
Management network:
 Gateway: <gatewayIP>
 IP Address/Mask: <nodeIP>/<subnet>
Cluster leader: yes

Re-enter config? (y/N): n
```

- Step 10** Repeat previous step to configure the initial information for the second and third nodes.

You do not need to wait for the first node configuration to complete, you can begin configuring the other two nodes simultaneously.

**Note**

You must provide the same password for all nodes or the cluster creation will fail.

The steps to deploy the second and third nodes are identical with the only exception being that you must indicate that they are not the **Cluster Leader**.

- Step 11** Wait for the initial bootstrap process to complete on all nodes.

After you provide and confirm management network information, the initial setup on the first node (`Cluster Leader`) configures the networking and brings up the UI, which you will use to add two other nodes and complete the cluster deployment.

```
Please wait for system to boot: [#####] 100%
System up, please wait for UI to be online.
```

**System UI online, please login to `https://<gatewayIP>72` to continue.**

- Step 12** Open your browser and navigate to `https://<node-mgmt-ip>` to open the GUI.

The rest of the configuration workflow takes place from one of the node's GUI. You can choose any one of the nodes you deployed to begin the bootstrap process and you do not need to log in to or configure the other two nodes directly.

Enter the password you entered in a previous step and click **Login**

**Step 13** After logging in, review the **Meet Nexus Dashboard** welcome screens, then proceed to the **Journey** screen. Click on **Go** under the **Cluster bringup** step to begin the bootstrap process.

**Step 14** Enter the requested information in the **Basic Information** page of the **Cluster Bringup** wizard.

- a) For **Cluster Name**, enter a name for this Nexus Dashboard cluster.

The cluster name must follow the [RFC-1123](#) requirements.

- b) For **Select the Nexus Dashboard Implementation type**, choose either **LAN** or **SAN** then click **Next**.

**Step 15** Enter the requested information in the **Configuration** page of the **Cluster Bringup** wizard.

- (Optional) If you want to enable IPv6 functionality for the cluster, put a check in the **Enable IPv6** checkbox.
- Click **+Add DNS provider** to add one or more DNS servers, enter the DNS provider IP address, then click the checkmark icon.
- (Optional) Click **+Add DNS search domain** to add a search domain, enter the DNS search domain IP address, then click the checkmark icon.
- (Optional) If you want to enable NTP server authentication, put a check in the **NTP Authentication** checkbox.
- If you enabled NTP authentication, click **+ Add Key**, enter the required information, and click the checkmark icon to save the information.
  - Key**—Enter the NTP authentication key, which is a cryptographic key that is used to authenticate the NTP traffic between the Nexus Dashboard and the NTP servers. You will define the NTP servers in the following step, and multiple NTP servers can use the same NTP authentication key.
  - ID**—Enter a key ID for the NTP host. Each NTP key must be assigned a unique key ID, which is used to identify the appropriate key to use when verifying the NTP packet.
  - Authentication Type**—Choose authentication type for the NTP key.
  - Put a check in the **Trusted** checkbox if you want this key to be trusted. Untrusted keys cannot be used for NTP authentication.

For the complete list of NTP authentication requirements and guidelines, see [General prerequisites and guidelines, on page 11](#).

If you want to enter additional NTP keys, click **+ Add Key** again and enter the information.

- If you enabled NTP authentication, click **+Add NTP Host Name/IP Address**, enter the required information, and click the checkmark icon to save the information.
  - NTP Host**—Enter an IP address; fully qualified domain names (FQDN) are not supported.
  - Key ID**—Enter the key ID of the NTP key you defined in the previous substep.  
If NTP authentication is disabled, this field is grayed out.
  - Put a check in the **Preferred** checkbox if you want this host to be preferred.

#### Note

If the node into which you are logged in is configured with only an IPv4 address, but you have checked **Enable IPv6** in a previous step and entered an IPv6 address for an NTP server, you will get the following validation error:

| NTP Host*                           | Key ID | Preferred |                                                                                                                                                                         |
|-------------------------------------|--------|-----------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 2001:420:28e:202a:5054:ff:fe6f:b3f6 |        | true      |   |

[+ Add NTP Host Name/IP Address](#)

 Could not validate one or more hosts Can not reach NTP on Management Network

This is because the node does not have an IPv6 address yet and is unable to connect to an IPv6 address of the NTP server. You will enter IPv6 address in the next step. In this case, enter the other required information as described in the following steps and click **Next** to proceed to the next page where you will enter IPv6 addresses for the nodes.

If you want to enter additional NTP servers, click **+Add NTP Host Name/IP Address** again and enter the information.

- g) For **Proxy Server**, enter the URL or IP address of a proxy server.

For clusters that do not have direct connectivity to Cisco cloud, we recommend configuring a proxy server to establish the connectivity. This allows you to mitigate risk from exposure to non-conformant hardware and software in your fabrics.

You can click **+Add Ignore Host** to enter one or more destination IP addresses for which traffic will skip using the proxy.

The proxy server must permit these URLs:

```
svc.intersight.com
svc-static1.intersight.com
svc-static1.ucs-connect.com
```

If you do not want to configure a proxy, click **Skip Proxy** then click **Confirm**.

- h) (Optional) If your proxy server requires authentication, put a check in the **Authentication required for Proxy** checkbox and enter the login credentials.
- i) (Optional) Expand the **Advanced Settings** category and change the settings if required.

Under advanced settings, you can configure these settings:

- **App Network**—The address space used by the application's services running in the Nexus Dashboard. Enter the IP address and netmask.
- **Service Network**—An internal network used by Nexus Dashboard and its processes. Enter the IP address and netmask.
- **App Network IPv6**—If you put a check in the **Enable IPv6** checkbox earlier, enter the IPv6 subnet for the app network.
- **Service Network IPv6**—If you put a check in the **Enable IPv6** checkbox earlier, enter the IPv6 subnet for the service network.

For more information about the application and service networks, see [General prerequisites and guidelines, on page 11](#).

- j) Click **Next**.

## Step 16

In the **Node Details** page, update the first node's information.

You have defined the Management network and IP address for the node into which you are currently logged in during the initial node configuration in earlier steps, but you must also enter the Data network information for the node before you can proceed with adding the other `primary` nodes and creating the cluster.

- a) For **Cluster Connectivity**, if your cluster is deployed in L3 mode, choose **BGP**. Otherwise, choose **L2**.

BGP configuration is required for the persistent IP addresses feature used by telemetry. This feature is described in more detail in the [BGP configuration and persistent IP addresses, on page 46](#) and [Nexus Dashboard persistent IP addresses, on page 40](#) sections.

### Note

You can enable BGP at this time or in the Nexus Dashboard GUI after the cluster is deployed. All remaining nodes need to configure BGP if it is configured. You must enable BGP now if the data network of nodes have different subnets.

- b) Click the **Edit** button next to the first node.

The node's **Serial Number**, **Management Network** information, and **Type** are automatically populated, but you must enter the other information.

- c) For **Name**, enter a name for the node.

The node's **Name** will be set as its hostname, so it must follow the [RFC-1123](#) requirements.

**Note**

If you need to change the name but the **Name** field is not editable, run the CIMC validation again to fix this issue.

- d) For **Type**, choose **Primary**.

The first nodes of the cluster must be set to **Primary**. You will add the secondary nodes in a later step if required for higher scale.

- e) In the **Data Network** area, enter the node's data network information.

Enter the data network IP address, netmask, and gateway. Optionally, you can also enter the VLAN ID for the network. Leave the VLAN ID field blank if your configuration does not require VLAN. If you chose **BGP** for **Cluster Connectivity**, enter the ASN.

If you enabled IPv6 functionality in a previous page, you must also enter the IPv6 address, netmask, and gateway.

**Note**

If you want to enter IPv6 information, you must do so during the cluster bootstrap process. To change the IP address configuration later, you would need to redeploy the cluster.

All nodes in the cluster must be configured with either only IPv4, only IPv6, or dual stack IPv4/IPv6.

- f) If you chose **BGP** for **Cluster Connectivity**, then in the **BGP peer details** area, enter the peer's IPv4 address and ASN.

You can click + **Add IPv4 BGP peer** to add additional peers.

If you enabled IPv6 functionality in a previous page, you must also enter the peer's IPv6 address and ASN.

- g) Click **Save** to save the changes.

**Step 17**

In the **Node Details** screen, click **Add Node** to add the second node to the cluster.

If you are deploying a single-node cluster, skip this step.

## Edit Node



### General

Name \*

nd-node1

Serial Number \*

E5998163D6F0

Type \*

Primary

### Management Network ⓘ

IPv4 Address/Mask \*

172.23.141.129/21

IPv4 Gateway \*

172.23.136.1

IPv6 Address/Mask

IPv6 Gateway

### Data Network ⓘ

IPv4 Address/Mask \*

172.31.140.68/21

IPv4 Gateway \*

172.31.136.1

IPv6 Address/Mask

IPv6 Gateway

VLAN ⓘ

Enable BGP ☐

Cancel

Save

- a) In the **Deployment Details** area, provide the **Management IP Address** and **Password** for the second node

You defined the management network information and the password during the initial node configuration steps.

- b) Click **Validate** to verify connectivity to the node.

The node's **Serial Number** and the **Management Network** information are automatically populated after connectivity is validated.

- c) Provide the **Name** for the node.  
d) From the **Type** dropdown, select `Primary`.

The first 3 nodes of the cluster must be set to `Primary`. You will add the secondary nodes in a later step if required for higher scale.

- e) In the **Data Network** area, provide the node's **Data Network** information.

You must provide the data network IP address, netmask, and gateway. Optionally, you can also provide the VLAN ID for the network. For most deployments, you can leave the VLAN ID field blank.

If you had enabled IPv6 functionality in a previous screen, you must also provide the IPv6 address, netmask, and gateway.

**Note**

If you want to provide IPv6 information, you must do it during cluster bootstrap process. To change IP configuration later, you would need to redeploy the cluster.

All nodes in the cluster must be configured with either only IPv4, only IPv6, or dual stack IPv4/IPv6.

- f) (Optional) If your cluster is deployed in L3 mode, **Enable BGP** for the data network.

BGP configuration is required for the persistent IP addresses feature. This feature is described in more detail in [BGP configuration and persistent IP addresses, on page 46](#) and the "Persistent IP Addresses" sections of the [Cisco Nexus Dashboard User Guide](#).

**Note**

You can enable BGP at this time or in the Nexus Dashboard GUI after the cluster is deployed.

If you choose to enable BGP, you must also provide the following information:

- **ASN** (BGP Autonomous System Number) of this node.  
You can configure the same ASN for all nodes or a different ASN per node.
- For IPv6-only, the **Router ID** of this node.  
The router ID must be an IPv4 address, for example `1.1.1.1`
- **BGP Peer Details**, which includes the peer's IPv4 or IPv6 address and peer's ASN.

- g) Click **Save** to save the changes.  
h) Repeat this step for the final (third) primary node of the cluster.

**Step 18** In the **Node Details** page, verify the information that you entered, then click **Next**.

**Step 19** Choose the **Deployment Mode** for the cluster.

- a) Click **Add Persistent Service IPs/Pools** to provide the required persistent IP addresses.

For more information about persistent IP addresses, see the [Nexus Dashboard persistent IP addresses, on page 40](#) section.

- b) Click **Next** to proceed.

**Step 20** In the **Summary** screen, review and verify the configuration information and click **Save** to build the cluster.

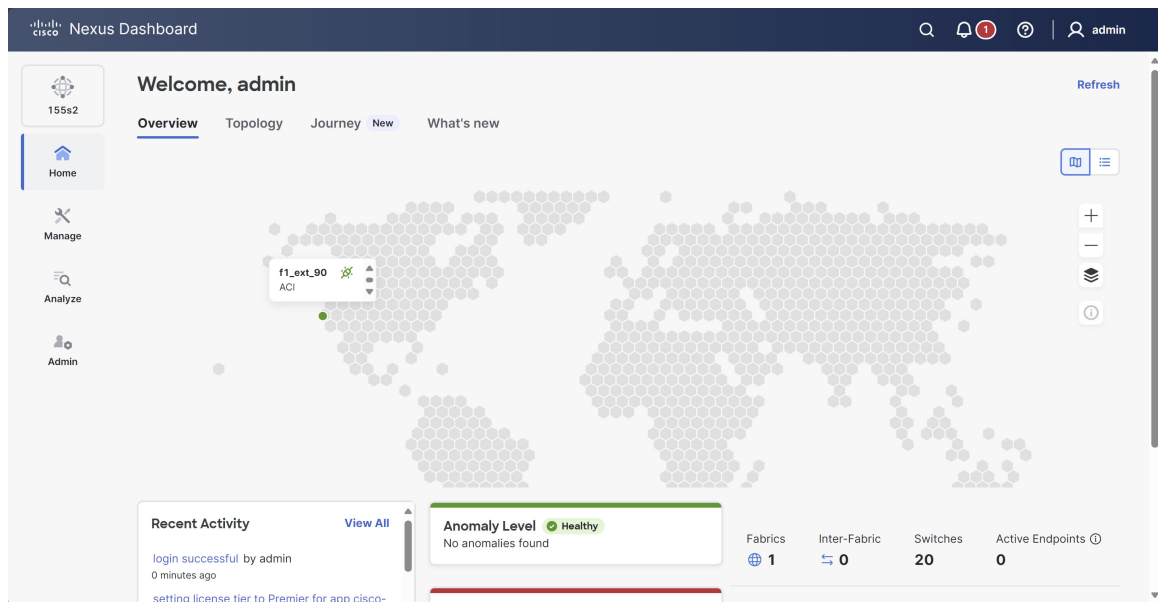
During the node bootstrap and cluster bring-up, the overall progress as well as each node's individual progress will be displayed in the UI. If you do not see the bootstrap progress advance, manually refresh the page in your browser to update the status.

It may take up to 30 minutes for the cluster to form and all the services to start. When cluster configuration is complete, the page will reload to the Nexus Dashboard GUI.

**Step 21** Verify that the cluster is healthy.

After the cluster becomes available, you can access it by browsing to any one of your nodes' management IP addresses. The default password for the `admin` user is the same as the `rescue-user` password you chose for the first node. During this time, the UI will display a banner at the top stating "Service Installation is in progress, Nexus Dashboard configuration tasks are currently disabled".

After all the cluster is deployed and all services are started, you can look at the **Anomaly Level** on the **Home > Overview** page to ensure the cluster is healthy:



Alternatively, you can log in to any one node using SSH as the `rescue-user` using the password you entered during node deployment and using the `acs health` command to see the status:

- While the cluster is converging, you may see the following output:

```
$ acs health
k8s install is in-progress

$ acs health
k8s services not in desired state - [...]

$ acs health
k8s: Etcd cluster is not ready
```

- When the cluster is up and running, the following output will be displayed:

```
$ acs health
All components are healthy
```

**Note**

In some situations, you might power cycle a node (power it off and then back on) and find it stuck in this stage:

```
deploy base system services
```

This is due to an issue with `etcd` on the node after a reboot of the physical Nexus Dashboard cluster.

To resolve the issue, enter the `acs reboot clean` command on the affected node.

**Step 22** (Optional) Connect your Cisco Nexus Dashboard cluster to Cisco Intersight for added visibility and benefits. Refer to [Working with Cisco Intersight](#) for detailed steps.

**Step 23** After you have deployed Nexus Dashboard, see the [collections page](#) for this release for configuration information.

---

**What to do next**

The next task is to create the fabrics and fabric groups. See the *Creating Fabrics and Fabric Groups* article for this release on the [Cisco Nexus Dashboard collections page](#).



## CHAPTER 9

# Deploying a Virtual Nexus Dashboard (vND) in Amazon Web Services (AWS)

- [About hosting a vND on the AWS public cloud, on page 203](#)
- [Prerequisites and guidelines for deploying the vNDs in Amazon Web Services, on page 205](#)
- [Prepare Amazon Web Services for the Nexus Dashboard cluster, on page 206](#)
- [Deploy a virtual Nexus Dashboard \(vND\) in Amazon Web Services \(AWS\), on page 208](#)
- [Node replacement \(RMA\), on page 214](#)

## About hosting a vND on the AWS public cloud

This feature allows you to run a virtual Nexus Dashboard (vND) on the AWS public cloud. The components to this solution are:

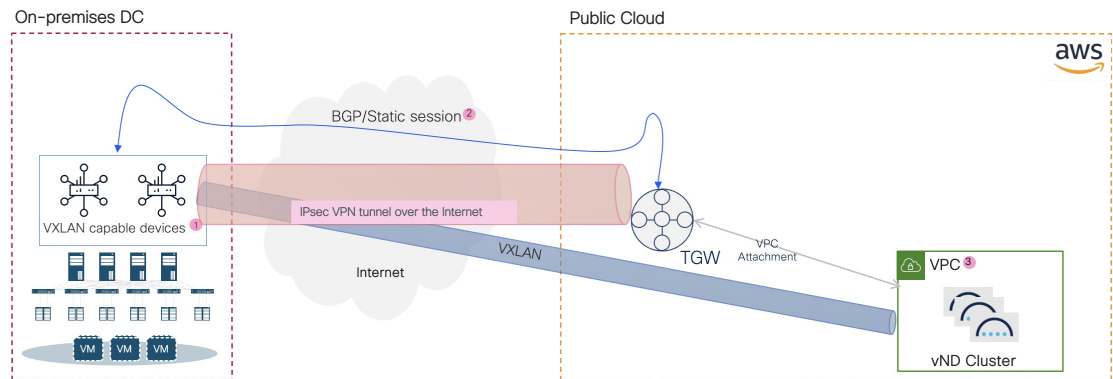
- Virtual Nexus Dashboard
- N9000 switch
- Two Catalyst 8000 series routers, or another type of device (such as the N9000 switch) that allows the Nexus Dashboard to terminate the VXLAN tunnel from the vND to the on-premises data center for persistent IP address (PIP) traffic
- AWS public cloud account

### Understanding how vNDs are deployed on AWS public cloud

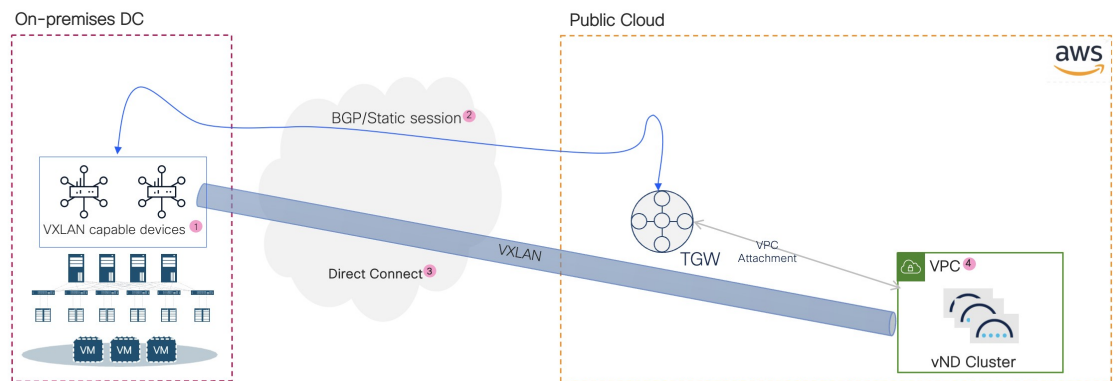
When you deploy the vNDs on the AWS public cloud, you do not have to do any manual bootstrapping; instead, the Nexus Dashboard bootstrapping performs the bootstrapping automatically for you. Once you go through the vND deployment process on AWS public cloud, a highly-available three-node cluster is created automatically for you across three availability zones (AZs) in a Virtual Private Cloud (VPC). In the Nexus Dashboard GUI, navigate to **Admin > System Status > Overview**, then view the information on the three-node cluster in the **Cluster nodes** area.

### Example topology

These figures show example topologies:



- ❶ On-premises device with VXLAN support for example, Catalyst 8000 series router, Nexus 9000
- ❷ BGP / static session for ND data IP reachability
- ❸ VPC in AWS terminology stands for virtual private cloud and is an equivalent of VRF



- ❶ On-premises device with VXLAN support for example, Catalyst 8000 series router, Nexus 9000
- ❷ BGP / static session for ND data IP reachability
- ❸ Direct connect takes the role of a private transport for carrying data
- ❹ VPC in AWS terminology stands for virtual private cloud and is an equivalent of VRF

Where:

- The connectivity between your on-premises data center and the AWS public cloud is achieved using either direct connect (recommended for production deployments) or an IPsec tunnel (for PoC or lab as additional overhead) between the two.
- The two on-premises routers could be one of the following:
  - If you are using direct connect, you can use two physical switches, such as the N9000 switches.
  - If you are using an IPsec tunnel, you can use a network appliance from the Cisco 8000 family, with VXLAN support because you will be terminating a VXLAN tunnel in this case.
- A transit gateway is used to create transit gateway attachments to connect to the VPC (the app VPC) that hosts the Nexus Dashboard nodes, as well as another transit gateway attachment, a VPN attachment or router, that is terminated on the transit gateway.



**Note** The Catalyst 8000V (C8000V) was launched as an evolution of the Cisco Cloud Services Router (CSR) 1000V. Throughout this documentation, C8000V will be used as an example of the VXLAN capable edge device. For more information, see [Release Notes for Cisco Catalyst 8000V Edge Software](#).

## Prerequisites and guidelines for deploying the vNDs in Amazon Web Services

Before you proceed with deploying the virtual Nexus Dashboards (vNDs) in Amazon Web Services (AWS):



**Note** The vND type that is supported for AWS is vND data only with 32vCPU, 128G RAM, 3TB SSD (GP3) and 10G of network throughput.

- This feature is supported on 3-node virtual clusters (data) on AWS with Nexus Dashboard as a single product (IPv4 only).
- Only LAN (NX-OS) and ACI fabrics are supported with this feature, and only with these enabled features:
  - Controller
  - Orchestration
  - Telemetry, with these restrictions:
    - Only out-of-band
    - Traffic analytics, but no flow telemetry
- This feature is supported on LAN and ACI fabrics only. It is not supported on IP Fabric for Media (IPFM) or SAN fabrics.



**Note** An AI fabric is considered a LAN fabric. Deploying this type of fabric in a vND in AWS is not restricted from a solution perspective. The basic AI fabric requirement is to not exceed latency between devices and the vND over 50ms.

- Secondary/worker nodes are not supported with this feature. Only three primary and one standby nodes are supported.
- Scale per cluster: 100 switches on a single three-node vND cluster
- You cannot change the IP address or the VNI assigned to the tunnel endpoints after you have deployed the vNDs in AWS.
- Before you can deploy your Nexus Dashboard vNDs, it is best to have your on-premises site ready, with the Catalyst 8000 series routers or a pair of N9000 switches already deployed, which allows you to

provide the necessary BDI and TEP IP addresses during the Nexus Dashboard vND deployment. If necessary, you can deploy the on-premises network appliances after you have deployed your Nexus Dashboard vNDs, but then you will have to ensure that you configure the on-premises devices with the same information that you provided during the Nexus Dashboard vND deployment. If you fail to provide the same configuration information in both places, you might have to re-install the Nexus Dashboard vNDs again.

- Verify that your on-premises VXLAN capable device is configured properly:

- Dataplane: Ingress Replication
- Control Plane: Flood and Learn

- Ensure that the AWS form factor supports your scale and services requirements.

Scale and services support and co-hosting vary based on the cluster form factor. You can use the [Nexus Dashboard Capacity Planning](#) tool to verify that the cloud form factor satisfies your deployment requirements.

- Review and complete the general prerequisites described in the [General prerequisites and guidelines, on page 11](#).
- Review and complete any additional prerequisites described in the *Release Notes* for the services you plan to deploy.
- Have appropriate access privileges for your AWS account.

You must be able to launch multiple instances of Elastic Compute Cloud (m5.xlarge) to host the Nexus Dashboard cluster.

- Ensure that the CPU family used for the Nexus Dashboard VMs supports AVX instruction set.
- Perform the [Prepare Amazon Web Services for the Nexus Dashboard cluster, on page 206](#) procedure.

## Prepare Amazon Web Services for the Nexus Dashboard cluster

Before you deploy the Nexus Dashboard vNDs in Amazon Web Services (AWS), follow these prerequisites to prepare AWS for your deployment:

- Familiarize yourself with AWS and how it works.
- (Optional) Establish a connection between AWS and your on-premises data center (ideally, a direct connection).
- Identify the region that you will use for the deployment of the Nexus Dashboard nodes.
- Have an existing VPC, or create a new VPC, that you will use for this deployment.
- Enable external access.

This is necessary for mapping Elastic IP addresses to the vND management interfaces and for accessing the GUI and SSH externally. You may or may not need to create and attach an Internet Gateway to the VPC to enable external access, depending on the connection method that you choose:

- **Option 1:** Connecting the management interface using the Ethernet Interface Processor (EIP), in which case you will need the Internet Gateway.

- **Option 2:** Use a private IP address, in which case you will not need the Internet Gateway.
- Update the security group to allow access from your public IP address or range for required services, such as:
  - HTTPS (TCP port 443): For accessing the Nexus Dashboard GUI
  - SSH (TCP port 22): For secure remote login to the vND nodes

This is necessary so that you can access the GUI and SSH into the Nexus Dashboard nodes.

- Create 6 subnets:
  - One set of subnets for management for each node (3) - minimum /28
  - One set of subnets for data for each node (3) - minimum /28

Subnets for management and data for a specific node must be in the same availability zone.

- Ensure that you have enough AWS Elastic IP addresses available for the vND deployment.

This installation requires 3 AWS Elastic IP addresses, 1 for each node, where each AWS Elastic IP address is used to access the management services, such as accessing the vND UI or SSH.

- Because the management subnets IP addresses will be mapped with AWS Elastic IP addresses as part of the deployment, those management subnets will need external access. Data subnets need to have reachability to the on-premises devices and on-premises Catalyst 8000 series routers (or other devices, such as N9000 switches) used for the termination of the VXLAN tunnel (used by persistent pods) from the vNDs.
- One /28 (such as 100.100.100.0/28) subnet that is not owned by AWS and comes from the on-premises data center (but is not already being used) that will be used by the PIP (persistent IP addresses), where 100.100.100.1 and 100.100.100.2 of that subnet must be the BDI IP addresses on the on-premises data center devices (Catalyst 8000 or N9000 switches) and the rest of the IP addresses are used by vND persistent pods (trap, telemetry collectors, and so on).
- On-premises devices should be able to reach these persistent IP addresses and Nexus Dashboard data IP addresses, and vice versa, for proper functioning.
- Create a security group with all the necessary IP addresses and ports so that the vND nodes can communicate with each other and form the cluster, and communicate externally and with the on-premises devices.
- Configure one EC2 key pair for deployment.



---

**Note** Complete this configuration as part of the prerequisites, even though EC2 key pair is not currently used and user/password is the only supported option at this time.

---

# Deploy a virtual Nexus Dashboard (vND) in Amazon Web Services (AWS)

This section describes how to deploy a virtual Nexus Dashboard (vND) in Amazon Web Services (AWS).

## Before you begin

- Ensure that you meet the requirements and guidelines described in [Prerequisites and guidelines for deploying the vNDs in Amazon Web Services, on page 205](#).

## Procedure

### Step 1

Subscribe to Cisco Nexus Dashboard product in AWS Marketplace.

- a) Log into your AWS account and navigate to the AWS Management Console.

The Management Console is available at <https://console.aws.amazon.com/>.

- b) Navigate to **Services > AWS Marketplace Subscriptions**.
- c) Click **Manage subscriptions**.
- d) Click **Discover products**.
- e) Search for **Cisco Nexus Dashboard - Cloud** and click the result.
- f) Click the **View Purchase** options and scroll down to click **Subscribe**.
- g) In the product page, click **View subscription**.
- h) In the **Manage subscriptions** page, locate the line for **Cisco Nexus Dashboard - Cloud**, then click **Launch** in that line.

### Step 2

Select software options and region.

- a) In the **Configure this software** page for **Cisco Nexus Dashboard - Cloud**, make the following choices:

- **Fulfillment option:** Leave the default **Nexus Dashboard - Cloud Deployment** choice as-is.
- **Software version:** Choose the latest 4.2.1 option available from the dropdown list.
- **Region:** Choose the appropriate region where the template will be deployed.

This must be the same region where you created your VPC.

- b) Click **Continue to Launch**.
- c) In the **Launch this software** page for **Cisco Nexus Dashboard - Cloud**, locate the **Choose Action** field and choose **Launch CloudFormation** from the dropdown list, then click **Launch**.

The **Create stack** page appears.

### Step 3

Complete the stack configuration.

- a) Leave the options in the **Create stack** page as-is.

#### Note

Do not make changes in the provided template. Only by using the smart default template configuration can you ensure a successful cluster formation.

- **Prerequisite - Prepare template:** Leave `Choose an existing template` option as-is.
- **Specify Template:** Leave `Amazon S3 URL` option as-is.
- **Amazon S3 URL:** Leave pre-configured URL entry as-is.

b) Click **Next** to continue.

The **Specify stack details** page appears.

#### Step 4

Specify the stack details.

- a) Provide the **Stack name**.
- b) Review the information provided in the **Parameters** area and make changes, if necessary.

For the most part, you can leave the pre-populated fields as-is based on the configurations that are part of this vND CFT.

- In the **Nexus Dashboard Cluster Name** field, enter the cluster name for the Nexus Dashboard cluster.
- In the **Fabric Deployment Mode** field, the default **LAN** option is the only supported option for the Nexus Dashboard 4.2.1 release.
- In the **VPC identifier** field, enter the VPC identifier.

The application VPC is automatically entered in this field. If you want to change the VPC in this field, choose another VPC under **VPC dashboard > Virtual private cloud > Your VPCs**.

- In the **Security Group Identifier** field, enter the security group identifier.

This is a pre-created security group that must allow ingress access for ports 22 and 443.

- In the **Instance Type** field, specify the EC2 instance type for the node instances.
- In the **AMI Identifier** field, specify the AWS AMI for the Nexus Dashboard.
- In the **Password** field, enter the admin password for the Nexus Dashboard node.

The admin password for the Nexus Dashboard node must contain at least 1 letter, number, and special character (@\$!%\*#?&) and must be between 8 and 64 characters in length.

- (Optional) In the **Key Pair Name** field, specify the name of an existing SSH key pair to enable SSH access to the Nexus Dashboard.

c) Enter the necessary information in the **DNS Configuration** area.

- In the **Primary DNS Server IP** field, enter the primary DNS server IP address.
- In the **Secondary DNS Server IP** field, enter the secondary DNS server IP address.
- In the **Search Domain Name** field, enter the search domain name.

d) (Optional) Enter the necessary information in the **Proxy Configuration** area.

- In the **Proxy Type** field, specify the proxy type (for example, HTTP or HTTPS).
- In the **Proxy URL** field, specify the full proxy URL, including the protocol and port (for example, `http://proxy.example.com:8080`).
- In the **Proxy Username** field, specify the proxy username, if authentication is required.

- In the **Proxy Password** field, specify the proxy password, if authentication is required.
- In the **Proxy Ignore Hosts IP** field, specify the proxy ignore hosts IP addresses.  
Only one entry is allowed in this field (for example, 192.168.10.101).

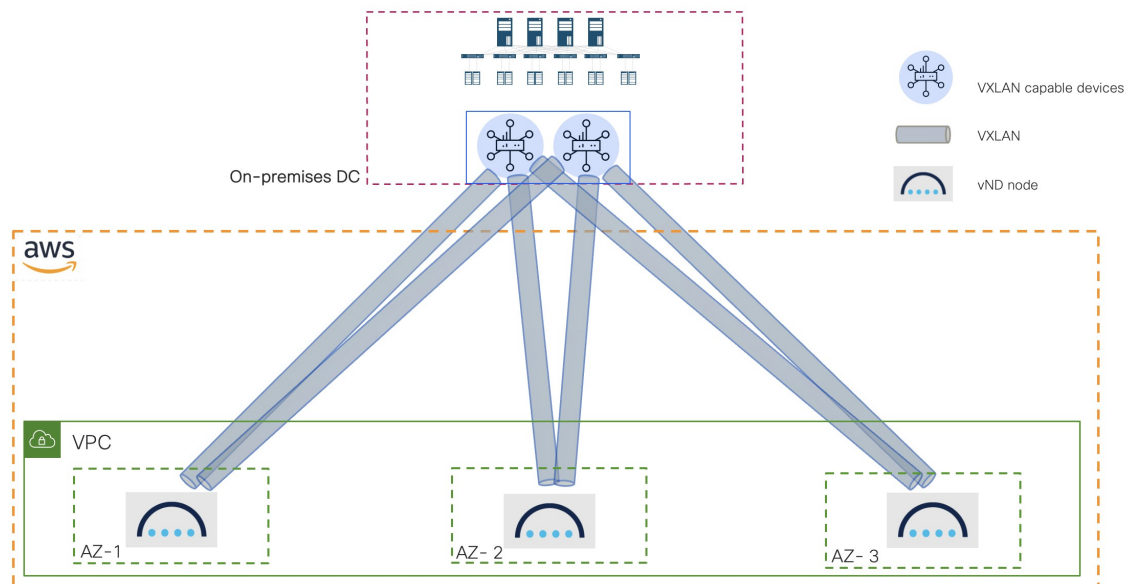
e) Enter the necessary information in the **NTP Configuration** area.

- In the **NTP Server Host** field, specify the NTP server host.
- In the **NTP Server Key Identifier** field, specify the NTP server key identifier.
- In the **NTP Server Preferred** field, choose `true` if the server is preferred.
- In the **NTP Key Identifier** field, specify the identifier for the NTP key.
- In the **NTP Key** field, specify the key for the NTP key.
- In the **NTP Key Authentication Type** field, specify the authentication type for the NTP key (for example, MDS or SHA1).
- In the **NTP Key Trusted** field, specify `true` or `false` for whether the NTP key is trusted.

f) Enter the necessary information in the **Cisco VXLAN Capable Device** area.

- In the **Device VXLAN Identifier (VNI)** field, enter a VNI value to be used for the VXLAN tunnel between the Cisco VXLAN capable device (if you are using an IPsec tunnel) or N9000 switches (if you are using direct connect) and the Nexus Dashboard nodes.

A single VNI value will be used for all the VXLAN tunnels between the Cisco VXLAN capable devices and the Nexus Dashboard nodes (the vNDs), as shown in this figure.



- In the **Device 1 Bridge Domain IP** and **Device 2 Bridge Domain IP** fields, enter the bridge domain IP addresses for both of the Cisco VXLAN capable devices.

The bridge domain IP addresses for the devices should come from the subnet that you provide in the **Private IP Subnet for Nexus Dashboard Pods** field. For example, if you enter 100.100.100.0/28 in the **Private IP**

**Subnet for Nexus Dashboard Pods** field, you might enter `100.100.100.1` and `100.100.100.2` as the bridge domain IP addresses for the devices.

- In the **Device 1 Tunnel Endpoint IP** and **Device 2 Tunnel Endpoint IP** fields, enter the tunnel endpoint IP addresses (the data IP addresses) for both of the Cisco VXLAN capable devices.
- In the **Private IP Subnet for Nexus Dashboard Pods** field, enter the private IP subnet to be used by the Nexus Dashboard pods.

The IP subnet size must be a /28, such as `100.100.100.0/28`.

**Note**

This procedures in this section do not deploy any Cisco VXLAN capable devices; they only ensure Nexus Dashboard has all the variables required to build connections with the edge devices.

- g) In the **Nexus Dashboard Node 1 Configuration**, **Nexus Dashboard Node 2 Configuration**, and **Nexus Dashboard Node 3 Configuration** areas, enter the necessary information for each of the vND nodes in the cluster:

- **ND Node x Hostname:** Enter the hostname for each Nexus Dashboard node.
- **ND Node x Management Subnet:** Enter the first management subnet each Nexus Dashboard node.
- **ND Node x Static Management IP:** Enter a static management IP address from the management subnet that you entered above for each Nexus Dashboard node.

Verify that the IP address that you enter in this field is not being used already.

- **ND Node x Management Subnet Netmask:** Enter the first management subnet netmask for each Nexus Dashboard node in the CIDR format (16-28).
- **ND Node x Management Subnet Gateway:** Enter the first management default gateway on the management subnet that you entered above for each Nexus Dashboard node.

This is typically the first address on the subnet.

- **ND Node x Data Subnet:** Enter the first data subnet each Nexus Dashboard node.
- **ND Node x Static Data IP:** Enter a static data IP address from the data subnet that you entered above for each Nexus Dashboard node.

Verify that the IP address that you enter in this field is not being used already.

- **ND Node x Data Subnet Netmask:** Enter the first data subnet netmask for each Nexus Dashboard node in the CIDR format (16-28).
- **ND Node x Data Subnet Gateway:** Enter the first data default gateway on the data subnet that you entered above for each Nexus Dashboard node.

This is typically the first address on the subnet.

- h) In the **Kubernetes Network Configuration (Optional)** area, enter the configuration information, if necessary:

- In the **Kubernetes Service Network** field, specify the network address for the Kubernetes service network.

The CIDR range is fixed to /16.

- In the **Kubernetes App Network** field, specify the network address for the Kubernetes app network.

The CIDR range is fixed to /16.

- i) Click **Next** to continue.

**Step 5** In the **Configure stack options** page, review and modify the information provided in this page, if necessary.

- a) Under **Stack failure options**, we recommend that you change the choice under **Behavior on provisioning failure** to **Preserve successfully provisioned resources**.
- b) Click **Next** when you have finished reviewing or modifying the information in the **Configure stack options** page.

**Step 6** In the **Review and create** page, verify the template configuration information, then click **Submit**.

**Step 7** Wait for the deployment to complete, then start the VMs.

You can view the status of the instance deployment in the **CloudFormation > Stacks** page, for example `CREATE_IN_PROGRESS`. You can click the refresh button in the top right corner of the page to update the status.

When the status for your stack changes to `CREATE_COMPLETE`, you can proceed to the next step.

**Step 8** In your stack under **CloudFormation > Stacks**, click the **Outputs** tab to view the public IP addresses for the three vNDs in the cluster.

**Note**

The CloudFormation template takes care of the connectivity within the cluster. Nodes should automatically form a cluster, if all the variables are filled in correctly.

**Step 9** Log into the Nexus Dashboard GUI using one of the public IP addresses listed in the previous step.

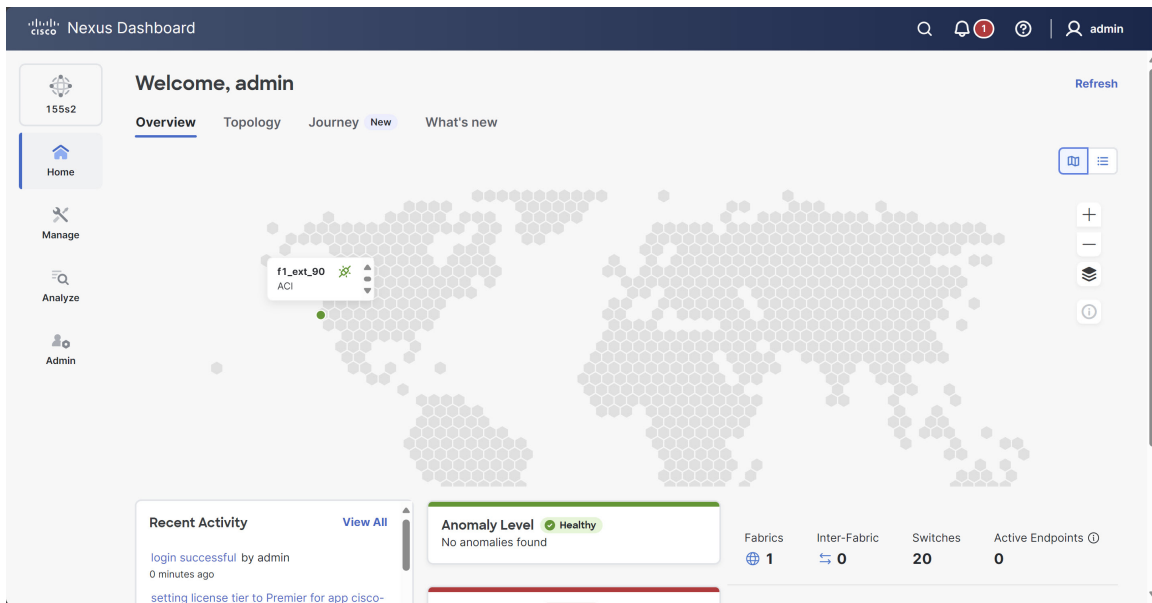
**Note**

You might have to wait for around 40 minutes after you have started the VMs before you can log into the Nexus Dashboard GUI using one of the public IP addresses.

**Step 10** Verify that the cluster is healthy.

After the cluster becomes available, you can access it by browsing to any one of your nodes' management IP addresses. The default password for the `admin` user is the same as the `rescue-user` password you chose for the first node. During this time, the UI will display a banner at the top stating "Service Installation is in progress, Nexus Dashboard configuration tasks are currently disabled".

After all the cluster is deployed and all services are started, you can look at the **Anomaly Level** on the **Home > Overview** page to ensure the cluster is healthy:



Alternatively, you can log in to any one node using SSH as the `rescue-user` using the password you entered during node deployment and using the `acs health` command to see the status:

- While the cluster is converging, you may see the following output:

```
$ acs health
k8s install is in-progress

$ acs health
k8s services not in desired state - [...]

$ acs health
k8s: Etcd cluster is not ready
```

- When the cluster is up and running, the following output will be displayed:

```
$ acs health
All components are healthy
```

#### Note

In some situations, you might power cycle a node (power it off and then back on) and find it stuck in this stage:

```
deploy base system services
```

This is due to an issue with `etcd` on the node after a reboot of the physical Nexus Dashboard cluster.

To resolve the issue, enter the `acs reboot clean` command on the affected node.

**Step 11** (Optional) Connect your Cisco Nexus Dashboard cluster to Cisco Intersight for added visibility and benefits. Refer to [Working with Cisco Intersight](#) for detailed steps.

**Step 12** After you have deployed Nexus Dashboard, see the [collections page](#) for this release for configuration information.

**What to do next**

The next task is to create the fabrics and fabric groups. See the *Creating Fabrics and Fabric Groups* article for this release on the [Cisco Nexus Dashboard collections page](#).

## Node replacement (RMA)

For AWS-based Nexus Dashboard deployments, if a node requires an RMA or replacement for any reason, a new replacement node must be deployed in AWS.

The factory-reset and node-recovery options available for on-premises physical Nexus Dashboard (pND) and virtual Nexus Dashboard (vND) deployments are not supported for AWS-based nodes. Therefore, the affected node cannot be recovered in place and must be redeployed as a new node.

Do not use the original three-node CloudFormation (CFT) template to replace an individual node. For single-node replacement, deploy the new node using the dedicated single-node CFT template provided below:

```
{
 "AWSTemplateFormatVersion": "2010-09-09",
 "Description": "This cloud formation template automates the creation and launching of
a Cisco Nexus Dashboard cluster (vND data 32vCPU/128GMem/10G Network/3T Disk) in multiple
Availability Zones (AZs) of AWS. Deployment requires a Application VPC to exist in the
account from which the cluster will be is launched. For detailed, description of this
template and usage instructions please refer to ND deployment guide.",
 "Metadata": {
 "AWS::CloudFormation::Interface": {
 "ParameterGroups": [
 {
 "Label": {
 "default": "Nexus Dashboard Basic Configuration"
 },
 "Parameters": [
 "NDClusterName",
 "FabricPersona",
 "VpcId",
 "SecurityGroupId",
 "InstanceType",
 "AMIID",
 "AdminPassword",
 "KeyPairName",
 "ExistingEIPAllocationId1",
 "ExistingEIPAddress1"
]
 },
 {
 "Label": {
 "default": "Nexus Dashboard Node Configuration"
 },
 "Parameters": [
 "hostNameND1",
 "SubnetNDMgmt1",
 "StaticNDMgmt1IP",
 "SubnetNDMgmt1Netmask",
 "SubnetNDMgmt1Gateway",
 "SubnetNDData1",
 "StaticNDData1IP",
 "SubnetNDData1Netmask",
 "SubnetNDData1Gateway"
]
 }
]
 }
 }
}
```

```

],
 "ParameterLabels": {
 "NDClusterName": { "default": "ND Cluster Name" },
 "VpcId": { "default": "VPC identifier" },
 "hostNameND1": { "default": "ND Node Hostname" },
 "SubnetNDMgmt1": { "default": "ND Node Management Subnet" },
 "StaticNDMgmt1IP": { "default": "ND Node Static Management IP" },
 "SubnetNDMgmt1Netmask": { "default": "ND Node Management Subnet Netmask"
 },
 "SubnetNDMgmt1Gateway": { "default": "ND Node Management Subnet Gateway"
 },
 "SubnetNDDat1": { "default": "ND Node Data Subnet" },
 "StaticNDDat1IP": { "default": "ND Node Static Data IP" },
 "SubnetNDDat1Netmask": { "default": "ND Node Data Subnet Netmask" },
 "SubnetNDDat1Gateway": { "default": "ND Node Data Subnet Gateway" },
 "SecurityGroupId": { "default": "Security Group identifier" },
 "InstanceType": { "default": "Instance Type" },
 "AMIId": { "default": "AMI identifier" },
 "AdminPassword": { "default": "Password" },
 "KeyPairName": { "default": "Key Pair Name" },
 "ExistingEIPAllocationId1": { "default": "ND Node 1 EIP Allocation ID." },
 "ExistingEIPAddress1": { "default": "ND Node 1 Elastic IP." }
 }
},
"Parameters": {
 "NDClusterName": {
 "Type": "String",
 "Description": "Nexus Dashboard Cluster name.",
 "MinLength": 1
 },
 "VpcId": {
 "Type": "String",
 "Description": "VpcId of your existing Virtual Private Cloud (VPC) to launch
ND cluster.",
 "ConstraintDescription": "Must be the VPC Id of an existing Virtual Private
Cloud."
 },
 "hostNameND1": {
 "Type": "String",
 "Description": "Specify the hostname for ND Node 1.",
 "Default": "AWSND1"
 },
 "SubnetNDMgmt1": {
 "Type": "AWS::EC2::Subnet::Id",
 "Description": "Select the 1st management subnet for the ND Node 1."
 },
 "StaticNDMgmt1IP": {
 "Type": "String",
 "Description": "Specify a static management IP address from the selected subnet.
Ensure not in-use already.",
 "AllowedPattern": "^((\\d{1,3})\\.\\d{1,3})\\.\\d{1,3})\\.\\d{1,3})$",
 "Default": "0.0.0.0"
 },
 "SubnetNDMgmt1Netmask": {
 "Type": "String",
 "Description": "Specify the 1st management subnet netmask in CIDR notation
(16-28).",
 "Default": "24",
 "AllowedPattern": "^(1[6-9]|2[0-8])$"
 },
 "SubnetNDMgmt1Gateway": {
 "Type": "String",
 "Description": "Specify the 1st management default gateway on the selected

```

```
management subnet. This is typically the first address on the subnet.",
 "AllowedPattern": "^((\\d{1,3})\\.){3}(\\d{1,3})$",
 "Default": "0.0.0.0"
},
"SubnetNDData1": {
 "Type": "AWS::EC2::Subnet::Id",
 "Description": "Select the 1st data subnet for the ND Node 1."
},
"StaticNDDatalIP": {
 "Type": "String",
 "Description": "Specify a static data IP address from the selected subnet. Ensure not in-use already.",
 "AllowedPattern": "^((\\d{1,3})\\.){3}(\\d{1,3})$",
 "Default": "0.0.0.0"
},
"SubnetNDDatalNetmask": {
 "Type": "String",
 "Description": "Specify the 1st data subnet netmask in CIDR notation (16-28).",
 "Default": "24",
 "AllowedPattern": "^(1[6-9]|2[0-8])$"
},
"SubnetNDDatalGateway": {
 "Type": "String",
 "Description": "Specify the 1st data default gateway on the selected management subnet. This is typically the first address on the subnet.",
 "AllowedPattern": "^((\\d{1,3})\\.){3}(\\d{1,3})$",
 "Default": "0.0.0.0"
},
"SecurityGroupId": {
 "Type": "AWS::EC2::SecurityGroup::Id",
 "Description": "Pre-created security group to be applied. Must allow ingress access for ports 22, 443. Check docs for other needed ports."
},
"InstanceType": {
 "Type": "String",
 "Description": "Specify EC2 instance type for the node instances.",
 "AllowedValues": ["m5.xlarge"],
 "Default": "m5.xlarge"
},
"AMIId": {
 "Type": "AWS::EC2::Image::Id",
 "Description": "Specify the AWS AMI for Nexus Dashboard."
},
"AdminPassword": {
 "Type": "String",
 "Description": "Admin user password for ND node (must contain atleast 1 letter, number and special char @$!%*#?& length: 8-64 chars).",
 "NoEcho": true,
 "MinLength": 8,
 "MaxLength": 64,
 "AllowedPattern": "^[^.*[A-Za-z])(?=.*\\d)(?=.*[@$!%*#?&])[A-Za-z\\d@$!%*#?&]{8,64}$"
},
"KeyPairName": {
 "Type": "AWS::EC2::KeyPair::KeyName",
 "Description": "Required - Specify Name of an existing SSH KeyPair to enable SSH access to ND."
},
"ExistingEIPAllocationId1": {
 "Type": "String",
 "Description": "Allocation ID of an existing Elastic IP for ND Node 1."
},
"ExistingEIPAddress1": {
```

```

 "Type": "String",
 "Description": "Public IP of the existing Elastic IP (e.g., 203.0.113.42) for
ND Node 1.",
 "Default": "0.0.0.0"
 }
},
"Resources": {
 "ND1": {
 "Type": "AWS::EC2::Instance",
 "Properties": {
 "InstanceType": { "Ref": "InstanceType" },
 "ImageId": { "Ref": "AMIID" },
 "KeyName": { "Ref": "KeyPairName" },
 "NetworkInterfaces": [
 {
 "DeviceIndex": 0,
 "NetworkInterfaceId": { "Ref": "ND1MgmtENI" }
 },
 {
 "DeviceIndex": 1,
 "NetworkInterfaceId": { "Ref": "ND1DataENI" }
 }
],
 "EbsOptimized": true,
 "BlockDeviceMappings": [
 {
 "DeviceName": "/dev/xvda",
 "Ebs": { "VolumeType": "gp3", "Iops": 5000 }
 },
 {
 "DeviceName": "/dev/xvdh",
 "Ebs": {
 "VolumeSize": 3072,
 "VolumeType": "gp3",
 "Iops": 5000
 }
 }
],
 "Tags": [{ "Key": "Name", "Value": "ND1-Master" }],
 "UserData": {
 "Fn::Base64": {
 "Fn::Sub": "\n \npassword\n": \"${AdminPassword}\",\n \n \"clusterName\n":
 \"${NDClusterName}\",\n \n \"oobNetwork\n\": {\n \n \n \"ipSubnet\n\":
 \"${StaticNDMgmtIP}/${SubnetNDMgmtNetmask}\",\n \n \n \"gateway\n\":
 \"${SubnetNDMgmtGateway}\",\n \n \n }\n\n"
 }
 }
 }
 },
 "ND1MgmtENI": {
 "Type": "AWS::EC2::NetworkInterface",
 "Properties": {
 "SubnetId": { "Ref": "SubnetNDMgmt1" },
 "PrivateIpAddresses": [
 {
 "PrivateIpAddress": { "Ref": "StaticNDMgmtIP" },
 "Primary": true
 }
],
 "GroupSet": [{ "Ref": "SecurityGroupId" }]
 }
 },
 "EIPAssoc1": {
 "Type": "AWS::EC2::EIPAssociation",

```

```

 "Properties": {
 "NetworkInterfaceId": { "Ref": "ND1MgmtENI" },
 "AllocationId": { "Ref": "ExistingEIPAllocationId1" }
 },
 "ND1DataENI": {
 "Type": "AWS::EC2::NetworkInterface",
 "Properties": {
 "SubnetId": { "Ref": "SubnetNDData1" },
 "PrivateIpAddresses": [
 {
 "PrivateIpAddress": { "Ref": "StaticNDData1IP" },
 "Primary": true
 }
],
 "GroupSet": [{ "Ref": "SecurityGroupId" }]
 }
 },
 "Outputs": {
 "ND1PublicIP": {
 "Description": "Public IP of ND Node 1.",
 "Value": { "Ref": "ExistingEIPAddress1" }
 }
 }
 }
}

```



## PART **III**

# Upgrading or Migrating to This Release

- [Upgrading a Nexus Dashboard 3.2.2 Cluster to This Release, on page 221](#)
- [Upgrading a Nexus Dashboard 4.1.1 Cluster to This Release, on page 239](#)
- [Migrating From DCNM to ND, on page 249](#)
- [Migrate from NDB to ND, on page 251](#)





## CHAPTER 10

# Upgrading a Nexus Dashboard 3.2.2 Cluster to This Release



### Note

The procedures in this chapter are applicable if you are upgrading from Nexus Dashboard release **3.2.2** to Nexus Dashboard release 4.3.1.

If you are upgrading from Nexus Dashboard release **4.1.1** to Nexus Dashboard release 4.3.1, follow the upgrade procedures provided in [Upgrading a Nexus Dashboard 4.1.1 Cluster to This Release, on page 239](#).

- [Prerequisites and guidelines for upgrading an existing Nexus Dashboard cluster, on page 221](#)
- [Nexus Dashboard pre-upgrade validation script, on page 226](#)
- [\(Optional\) Convert separate NDFC and NDI clusters to a single co-hosted Nexus Dashboard cluster, on page 227](#)
- [Supported upgrade paths, on page 228](#)
- [Upgrade Nexus Dashboard, on page 230](#)
- [Post-upgrade information and tasks, on page 233](#)
- [Troubleshooting upgrades, on page 237](#)

## Prerequisites and guidelines for upgrading an existing Nexus Dashboard cluster

Before you upgrade your existing Nexus Dashboard cluster:

- Ensure that you have read the target release's [Release Notes](#) for any changes in behavior, guidelines, and issues that may affect your upgrade.
- Before upgrading to Nexus Dashboard release 4.3.1:
  - Make sure your NTP and DNS services are configured. At least one NTP and DNS are required for the system to upgrade successfully.
  - For LAN deployments, verify that the management network and data network are in different subnets. The upgrade will fail if the management network and data network are not in different subnets.

Note that this restriction does not apply for SAN deployments (management network and data network can be in same subnet for SAN deployments).

- We highly recommend that you use the Nexus Dashboard Pre Upgrade Validation script before performing any Nexus Dashboard upgrades. See [Nexus Dashboard pre-upgrade validation script, on page 226](#) for more information.

- Verify that the `acs health` is healthy.

1. Access the Nexus Dashboard using `ssh -l rescue-user {management-ip-of-nd}`.
2. Issue the `acs health` command.

The output from the `acs health` command should show that all components are healthy:

```
rescue-user@node1:~$ acs health
=====
Status
=====
All components are healthy
```

- Review these important guidelines about how backup and restore processes affect your upgrade:
  - Ensure that you perform a backup of your Nexus Dashboard cluster before upgrading and you store the backup file in a safe place. To perform a backup, refer to [Unified Backup and Restore for Nexus Dashboard and Services](#). Note that you will not be able to restore this backup directly to a Nexus Dashboard cluster running the 4.3.1 release.
  - An upgrade will not proceed if the most recent backup had a failure. Make sure you have a successful backup before proceeding with the upgrade. If you are unable to perform a successful backup and cannot upgrade, contact [Cisco Technical Assistance Center \(TAC\)](#) for support.
  - You will not be able to perform an upgrade if a restore has failed. After the restore has completed, click **View History** to navigate to the **History** area in the **Backup and Restore** page, as described in the "Restore Nexus Dashboard configurations" section in the [Unified Backup and Restore for Nexus Dashboard and Services](#).

The page should display **Success** in the **Status** column for the restore process. If you see any value other than **Success** in the **Status** column, redeploy the cluster and get a **Success** value for the restore process before attempting the upgrade again.

- If you have the **Enable NX-API certificate verification** field enabled under **Admin > Certificate Management > Fabric Certificates** and you upgrade from Nexus Dashboard release 3.2.2 to 4.3.1, those certificates might be lost after the upgrade. We recommend that you run the Nexus Dashboard Pre Upgrade Validation script before performing any Nexus Dashboard upgrades, which can proactively identify conditions that might lead to certificate loss or misconfiguration post-upgrade. See [Nexus Dashboard pre-upgrade validation script, on page 226](#) for more information.
- Beginning with Nexus Dashboard release 4.3.1, the App node with 1.5TB disk cluster type is no longer supported as a greenfield deployment.
- If NDI performs telemetry for remotely created NDFC clusters, or if you use multi-cluster connectivity across multiple Nexus Dashboard clusters, upgrade all clusters to release 4.3.1. Do not run a mix of release 3.2x and 4.3.1 clusters with multi-cluster connectivity. Upgrade co-located NDI and NDFC clusters to release 4.3.1 as well. Both clusters must run release 4.3.1 before you can use the new Connected TAC functionality.

- If you are upgrading a physical Nexus Dashboard cluster, ensure that the nodes have the minimum supported CIMC version for the target Nexus Dashboard release.

Supported CIMC versions are listed in the [Nexus Dashboard Release Notes](#) for the target release.

The CIMC upgrade is described in detail in the "Troubleshooting" article in the Nexus Dashboard [documentation library](#).

- If you are upgrading a virtual Nexus Dashboard cluster, Nexus Dashboard will enforce these checks:
  - A check of the HDD latency to verify that it is <30ms. If the HDD has a higher latency, the upgrade will fail.
  - A check of the network latency to verify that it is <50ms within cluster nodes. If the network has a higher latency, the upgrade will fail.
- If you are upgrading a virtual Nexus Dashboard cluster deployed in VMware ESX, ensure that the ESX version is still supported by the target release.

This release supports VMware ESXi 7.0, 7.0.1, 7.0.2, 7.0.3, 8.0, 8.0.2, 8.0.3.



**Note** If you need to upgrade the ESX server, you must do that before upgrading your Nexus Dashboard. ESX upgrades are outside the scope of this document, but in short:

1. Upgrade one of the ESX hosts as you typically would with your existing Nexus Dashboard node VM running.
2. After the host is upgraded, ensure that the Nexus Dashboard cluster is still operational and healthy.
3. Repeat the upgrade on the other ESX hosts one at a time.
4. After all ESX hosts are upgraded and the existing Nexus Dashboard cluster is healthy, proceed with upgrading your Nexus Dashboard to the target release as described in this document.

- Ensure that your current Nexus Dashboard cluster is healthy.

You can check the system status on the **Overview** page of the Nexus Dashboard's **Admin Console** or by logging in to one of the nodes as `rescue-user` and ensuring that the `acs health` command returns `All components are healthy`.

- Nexus Dashboard does not support platform downgrades.

If you want to downgrade to an earlier release, you will need to deploy a new cluster.

- If you have a user who only has the `Dashboard User` (`app-user`) user role in Nexus Dashboard release 3.2.x, after the upgrade to Nexus Dashboard release 4.3.1, delete the user with the `Dashboard User` user role or use the `Observer` role instead for that user in Nexus Dashboard release 4.3.1.
- The number of persistent IP addresses and how they are mapped out has changed from previous releases to Nexus Dashboard release 4.3.1. See [Nexus Dashboard persistent IP addresses, on page 40](#) to understand the number of persistent IP addresses that were needed in previous releases, and how you will have to

make certain updates in the number of persistent IP addresses that you will need before you can upgrade to Nexus Dashboard release 4.3.1.

- If you have a Nexus Dashboard (ND) in any persona with any services, specifically Nexus Dashboard Fabric Controller (NDFC) and Nexus Dashboard Insights (NDI), and you have upgraded from previous releases (such as ND 2.2.x and below) to ND 3.2.2, be aware of the following important information regarding Elasticsearch (ES) indices size:

If the cluster was deployed on ND 2.2.x or earlier and was upgraded since then, the upgrade from ND 3.2.2 to ND 4.3.1 may be indexing the time series database. The process will prompt you to contact [Cisco Technical Assistance Center \(TAC\)](#) if the sizes exceed the sizes of the ability to reindex on this cluster. If the process fails to reindex, you can use the `acs recover` command that's provided in the UI to proceed with the upgrade after the failure.

- If you have Nexus Dashboard Orchestrator as part of your pre-ND release 4.3.1 cluster deployed using the same or different data subnets, before upgrading to Nexus Dashboard release 4.3.1:
  - You must add persistent IP addresses. See [Nexus Dashboard persistent IP addresses, on page 40](#).
  - Additionally, if your pre-ND release 4.3.1 cluster nodes uses different data subnets, you must also add BGP configurations on all the nodes.

The validation of the image during the upgrade will fail if you do not have these items configured before upgrading.

- Nexus Dashboard release 4.1x introduced the **Access** switch role, which was not available in Nexus Dashboard releases 3.1x or 3.2x.
  - **3.x behavior:** Telemetry exporter uses Loopback as source interface
  - **4.x behavior:** Telemetry exporter uses SVI as source interface

This change is triggered automatically after an upgrade.

Flow Telemetry exporters cannot accept a modification of the source-interface once its created. Any attempt to change it (Loopback → SVI) will fail on the switch.

When upgrading from Nexus Dashboard releases 3.1x or 3.2x to 4.x, redeploying a Flow Telemetry configuration might fail due to the source-interface mismatch (Loopback → SVI).

### Auto reconcile configuration drifts on upgrade

When upgrading to ND release 4.3.1 from any pre-4.2.(1) NDO release, you must perform a multi-step upgrade:

1. First upgrade the pre-4.2(1) NDO release to ND 3.0.1i, then to ND release 3.2.2, as described in the ["Supported Upgrade Paths"](#) section in *Cisco Nexus Dashboard and Services Deployment and Upgrade Guide, Release 3.2.2*.
2. Then upgrade from ND release 3.2.2 to ND release 4.3.1, as described in this chapter.

When upgrading a pre-4.2(1) NDO release to ND release 3.2.2, you may observe some configuration drifts in application templates. These drifts occur because NDO release 4.2(1) and later support managing new application template object properties that were not managed in earlier versions. For a list of new properties managed by NDO 4.2(1), refer to the [Nexus Dashboard Orchestrator 4.2\(1\) release notes](#).

ND release 4.3.1 supports automatic reconciliation of configuration drifts as part of the upgrade process. Nexus Dashboard will check whether the templates are in sync with the fabrics and, if necessary, will import fabric values into Nexus Dashboard to automatically resolve the drifts.

When following the multi-step upgrade path from a pre-4.2(1) NDO version, we recommend that you ignore any drifts detected in ND release 3.2.2 and continue the upgrade to ND release 4.3.1.

Any drifts remaining after upgrading to ND release 4.3.1 will need to be resolved manually. For example, a drift that cannot be auto-resolved occurs when a configured object is part of a stretched template across multiple fabrics, and the template-level property has different values configured on different fabrics.

### (Optional) Convert a 6-node pND cluster to a 3-node pND cluster

If you want to convert a 6-node pND cluster to a 3-node pND as part of the upgrade process (for example, if you have a 6-node pND controller deployment in Nexus Dashboard release 3.2.2, which isn't supported in Nexus Dashboard release 4.1.1), use these procedures to make this conversion in Nexus Dashboard release 3.2.2, before upgrading to Nexus Dashboard release 4.1.1.

Note that a 6-node pND **telemetry**-only deployment is supported in Nexus Dashboard release 4.1.1 so this optional procedure is not needed in this case.

1. Perform a backup of the 6-node pND cluster in Nexus Dashboard release 3.2.2.  
See [Unified Backup and Restore for Nexus Dashboard and Services](#) for more information.
2. Ensure that the primary nodes and the cluster are healthy, then delete the secondary nodes one by one.
  - a. In your Nexus Dashboard running on release 3.2.2, navigate to **Manage > Nodes**.
  - b. Select the secondary node that you want to delete.
  - c. From the **Actions** menu, choose **Delete** to delete the node.
  - d. Navigate to **Overview > Platform View** and wait until the cluster health status shows **GREEN** after each secondary node deletion. Alternatively, verify the cluster health by running the `acs health` command. Do not delete another secondary node until the cluster health is healthy.



---

**Note** Kafka may take up to 30 minutes to stabilize after deleting a node.

---

- e. Once you see the cluster health status as **GREEN** after you deleted a secondary node, proceed to delete the next secondary node.  
  
Continue the process of deleting each secondary node one-by-one, waiting for the cluster status to turn **GREEN** before proceeding to the deletion of the next secondary node.
3. When you have deleted all the secondary nodes, verify that the cluster is healthy, then perform a new backup in Nexus Dashboard release 3.2.2.  
See [Unified Backup and Restore for Nexus Dashboard and Services](#) for more information.
4. Upgrade the cluster from Nexus Dashboard release 3.2.2 to release 4.3.1 using the procedures in this chapter.

# Nexus Dashboard pre-upgrade validation script

We highly recommend that you use the Nexus Dashboard pre-upgrade validation script before performing any Nexus Dashboard upgrades. The Nexus Dashboard pre-upgrade validation script performs various checks for known issues that have been identified to affect the success of a Nexus Dashboard upgrade. The script is continuously updated and maintained in an effort to mitigate any new upgrade-related issues that are detected in the field.

- Prior to Nexus Dashboard release 4.3.x, the Nexus Dashboard pre-upgrade validation script was a standalone Python script that you would [download](#) and run from an external host over SSH to perform health checks across nodes.
- Beginning with Nexus Dashboard release 4.3.x, the Nexus Dashboard pre-upgrade validation script is now a signed plugin process, where you download the Cisco-signed `.ndp` (Nexus Dashboard plugin) file, which Nexus Dashboard validates and runs locally with elevated privileges.

## Guidelines and limitations: Pre-upgrade validation script commands

These are the guidelines and limitations for the pre-upgrade validation script commands:

- Only one plugin can be installed at a time. If you want to download the Nexus Dashboard pre-upgrade validation script plugin file and you already have a plugin file installed, you must remove the installed plugin first, then you can download and run another Nexus Dashboard pre-upgrade validation script plugin file. See [Pre-upgrade validation script commands, on page 226](#) for more information.

## Pre-upgrade validation script commands

The pre-upgrade validation script commands are `acs` CLI commands, similar to the Nexus Dashboard `acs` CLI commands that are described in the [Cisco Nexus Dashboard Troubleshooting](#) article. Log into any of the nodes in the Nexus Dashboard cluster to run these pre-upgrade validation script commands `acs` CLI commands.

These are the commands that are available for the new Nexus Dashboard `.ndp` version of the pre-upgrade validation script:

- To download the Nexus Dashboard pre-upgrade validation script plugin file, enter:

```
acs plugin download <url>
```

where `<url>` is the plugin download URL. This command downloads, verifies, and installs the plugin file.

The supported methods are positional arguments, and the supported download formats are:

- **HTTP download:** `http://<server>/<path>`. For example:

```
acs plugin download http://192.168.1.100/my-plugin.ndp
```

- **HTTPS download:** `https://<server>/<path>`. For example:

```
acs plugin download https://server.example.com/plugins/my-plugin.ndp
```

- **SFTP from remote secure storage:** `sftp://<server>:/<path>`. For example:

```
acs plugin download sftp://sftp-server.example.com:/export/plugins/my-plugin.ndp
```

- **SCP from remote secure storage:** `scp://<server>:/<path>`. For example:

```
acs plugin download scp://scp-server.example.com:/export/plugins/my-plugin.ndp
```

**Note**

- In order to use the SCP and SFTP download formats, you must first configure the remote storage. Refer to [Create a remote storage location](#) for those procedures and choose the **SFTP/SCP Server** option in the **Remote Storage Location Type** field.
- SCP (scp://) also supports SFTP for backwards compatibility.

If you already have a plugin file installed, you will see the message `plugin already installed`, run `'acs plugin delete'` first. Since you can only have one installed plugin file at a time, determine if you want to replace the existing plugin file:

- Enter `acs plugin list` to determine the version for the plugin that's currently installed. You can choose not to install the new plugin file if you see that the version is the same.
- If you want to replace the existing plugin file, remove the installed plugin file using the `acs plugin delete` command as described below, then attempt to download the new plugin file again.
- To run the Nexus Dashboard pre-upgrade validation script plugin file, enter:

```
acs plugin run
```

If you do not have a plugin installed, the output returns a `no plugin installed` message.

- To show plugin file information, enter:

```
acs plugin list
```

- If you have a plugin installed, the output returns the name, version, and `signed_by` information.
- If you do not have a plugin installed, the output returns a `no plugin installed` message.

- To remove an installed plugin file, enter:

```
acs plugin delete
```

Since you can only have one installed plugin file at a time, use this command to remove an already installed plugin file if you are attempting to download a new plugin file.

## (Optional) Convert separate NDFC and NDI clusters to a single co-hosted Nexus Dashboard cluster

If you want to convert separate NDFC and NDI clusters in Nexus Dashboard release 3.x to a single co-hosted Nexus Dashboard cluster in release 4.x as part of the upgrade process, use these procedures before upgrading.

For example, if you have a 5-node vND cluster running NDFC and a separate 6-node vND cluster running NDI in Nexus Dashboard release 3.2.2, and you want to merge them into a single Nexus Dashboard release 4.x cluster, complete the conversion steps in Nexus Dashboard release 3.2.2 before upgrading.

## Procedure

- 
- Step 1** On the NDI cluster, disable telemetry completely.
- Remove the NDFC fabrics from the NDI cluster.
  - If the NDFC and NDI clusters are federated (connected through multi-cluster connectivity), disconnect the clusters and remove the federation configuration.
- See [Nexus Dashboard Infrastructure Management, Release 3.2.x](#) for more information.
- Step 2** Upgrade only the Nexus Dashboard cluster running NDFC from release 3.2.2 to release 4.x using the procedures in this chapter.
- Step 3** After the upgrade is complete, verify that all NDFC functions and features are working as expected.
- Step 4** Perform a backup of the upgraded Nexus Dashboard release 4.x cluster.
- See [Backing Up and Restoring Your Nexus Dashboard](#) for more information.
- Step 5** Enable telemetry from the fabric settings.
- For a 3-node or larger pND cluster, enable telemetry directly from the fabric settings after verifying the NDFC upgrade.
  - If the upgraded deployment is a 3-node vND app cluster, create a new 3-node data vND or 6-node vND (3 data + 3 app) cluster or 3- or 6-pND cluster, then restore the Nexus Dashboard release 4.x backup taken in the previous step.
  - After the restore is complete, verify all NDFC functions and features, then enable telemetry from the fabric settings.

### Note

Do not upgrade the separate NDI cluster to Nexus Dashboard release 4.x as part of this conversion. Telemetry should be re-enabled from the upgraded NDFC cluster after the co-hosted deployment is ready.

---

## Supported upgrade paths

As described in [Nexus Dashboard deployment overview, on page 5](#), in earlier releases, Nexus Dashboard shipped with only the platform software and no services included, which you would then download, install, and enable separately after the initial platform deployment. In addition, Nexus Dashboard release 3.1.1 introduced a tighter coupling between the Nexus Dashboard and individual services with only a single version of each service compatible with each version of the platform. As a result, as long as you were on the minimum required version of the Nexus Dashboard software, you could upgrade both the platform and all currently enabled services directly to Nexus Dashboard release 3.1x and 3.2x.

Beginning with Nexus Dashboard release 4.1.1, the platform and the individual services have been unified into a single product, which means that you no longer deploy, configure, or upgrade the services separately.

The following table provides a few example scenarios for specific deployment combinations:

Table 28:

| Current Nexus Dashboard Release | Compatible Services<br>(depending on form factor and cluster size, you may have one or more of these services currently enabled) | Upgrade Workflow                                                                                                                                                                                                                                                                                                                                                                                                                                                   |
|---------------------------------|----------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 4.1.1                           | N/A                                                                                                                              | Upgrade directly to release 4.3.1 as described in <a href="#">Upgrading a Nexus Dashboard 4.1.1 Cluster to This Release, on page 239</a> .                                                                                                                                                                                                                                                                                                                         |
| 3.2.2                           | Fabric Controller: 12.2(3)<br>Orchestrator: 4.4(2)<br>Insights: 6.5(2)                                                           | Upgrade directly to release 4.3.1 as described in the following section.<br><br>All services are unified under a single Nexus Dashboard product in release 4.3.1.                                                                                                                                                                                                                                                                                                  |
| 3.2.1                           | Fabric Controller: 12.2(2)<br>Orchestrator: 4.4(1)<br>Insights: 6.5(1)                                                           | <ol style="list-style-type: none"> <li>1. Upgrade the Nexus Dashboard platform to release 3.2.2 as described in <a href="#">Nexus Dashboard Deployment Guide, Release 3.2.x</a><br/><br/>All services will be automatically upgraded along with the platform.</li> <li>2. Upgrade from release 3.2.2 to release 4.3.1 as described in the following section.<br/><br/>All services are unified under a single Nexus Dashboard product in release 4.3.1.</li> </ol> |
| 3.1.1                           | Fabric Controller: 12.2(1)<br>Orchestrator: 4.3(x)<br>Insights: 6.4(1)                                                           | <ol style="list-style-type: none"> <li>1. Upgrade the Nexus Dashboard platform to release 3.2.2 as described in <a href="#">Nexus Dashboard Deployment Guide, Release 3.2.x</a><br/><br/>All services will be automatically upgraded along with the platform.</li> <li>2. Upgrade from release 3.2.2 to release 4.3.1 as described in the following section.<br/><br/>All services are unified under a single Nexus Dashboard product in release 4.3.1.</li> </ol> |
| 3.0.1                           | Fabric Controller: 12.1(3)<br>Orchestrator: 4.2(x)<br>Insights: 6.3(1)                                                           | <ol style="list-style-type: none"> <li>1. Upgrade the Nexus Dashboard platform to release 3.2.2 as described in <a href="#">Nexus Dashboard Deployment Guide, Release 3.2.x</a><br/><br/>All services will be automatically upgraded along with the platform.</li> <li>2. Upgrade from release 3.2.2 to release 4.3.1 as described in the following section.<br/><br/>All services are unified under a single Nexus Dashboard product in release 4.3.1.</li> </ol> |

| Current Nexus Dashboard Release | Compatible Services<br>(depending on form factor and cluster size, you may have one or more of these services currently enabled) | Upgrade Workflow                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                            |
|---------------------------------|----------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 2.3.2 and earlier               | Fabric Controller: 12.1(2) or earlier<br>Orchestrator: 4.1(x) or earlier<br>Insights: 6.2(x) or earlier                          | <ol style="list-style-type: none"> <li>1. Upgrade the Nexus Dashboard platform to release 3.1.1 as described in <a href="#">Nexus Dashboard Deployment Guide, Release 3.1.x</a>.<br/>All services will be automatically upgraded along with the platform.</li> <li>2. Upgrade from release 3.1.1 to release 3.2.2 as described in <a href="#">Nexus Dashboard Deployment Guide, Release 3.2.x</a>.<br/>All services will be automatically upgraded along with the platform.</li> <li>3. Upgrade from release 3.2.2 to release 4.3.1 as described in the following section.<br/>All services are unified under a single Nexus Dashboard product in release 4.3.1.</li> </ol> |

## Upgrade Nexus Dashboard

This section describes how to upgrade an existing Nexus Dashboard cluster if you are upgrading from Nexus Dashboard release 3.2.2 to Nexus Dashboard release 4.3.1.



**Note** If you are upgrading from Nexus Dashboard release 4.1.1 to Nexus Dashboard release 4.3.1, follow the upgrade instructions provided in [Upgrading a Nexus Dashboard 4.1.1 Cluster to This Release, on page 239](#).

As described in [Supported upgrade paths, on page 228](#), you must be on Nexus Dashboard release 3.2.2 to upgrade directly to Nexus Dashboard release 4.3.1. Note these important points on these releases:

- **Nexus Dashboard release 3.2.x:** As described in [Nexus Dashboard deployment overview, on page 5](#), for these Nexus Dashboard releases, Nexus Dashboard was available as one product, and individual services (such as Nexus Dashboard Insights, Orchestrator, and Fabric Controller) were available as individual products, separate from Nexus Dashboard.
- **Nexus Dashboard release 4.1.1:** This is the first Nexus Dashboard release where Nexus Dashboard and the individual services listed above are packaged together as a single, unified product. Nexus Dashboard release 4.3.1 is the next unified Nexus Dashboard release following the 4.1.1 release.

This means that your upgrade process will go through these phases:

1. You will begin the upgrade process in Nexus Dashboard release 3.2.2, where Nexus Dashboard and the individual services are separate products.
2. You will then upgrade to Nexus Dashboard 4.3.1, where you will end the upgrade process with Nexus Dashboard and the individual services packaged together as a single, unified product.

### Before you begin

Ensure that you have completed the prerequisites described in [Prerequisites and guidelines for upgrading an existing Nexus Dashboard cluster, on page 221](#)

### Procedure

#### Step 1

In your Nexus Dashboard release 3.2.2 system, download the Nexus Dashboard 4.3.1 image.

- a) Browse to the Software Download page.

<https://software.cisco.com/download/home/286327743/type/286328258>

- b) Choose the Nexus Dashboard 4.3.1 release to download.

- c) Download the Nexus Dashboard image for the 4.3.1 release.

#### Note

- The upgrade process is the same for all Nexus Dashboard form factors and uses the Nexus Dashboard ISO image (nd-dk9.<version>.iso). In other words, even if you used the virtual form factors (such as the ESX .ova) or a cloud provider's marketplace for initial cluster deployment, you must still use the .iso image for upgrades.
- If the image fails to download completely, verify the network connectivity between the Nexus Dashboard and the image server. Check the proxy configuration under **Admin > System Settings > General > Proxy configuration**.

- d) (Optional) Host the image on a web server in your environment.

#### Note

We recommend hosting the image on a server in your environment. When you upload the image to your Nexus Dashboard cluster, you will have an option to provide a direct URL to the image, which can significantly speed up the process.

#### Step 2

Log in to your current Nexus Dashboard's **Admin Console** as an `Administrator` user.

#### Step 3

Delete any older, non-active upgrade images from your cluster.

If this is the first time you're upgrading your cluster, you can skip this step.

- a) Navigate to **Manage > Software Management**.
- b) Click the trash icon on an upgrade image's tile to delete any older, non-active upgrade **Images**.
- c) Repeat this step for all older, non-active upgrade images.

#### Step 4

Upload the new image to the cluster.

- a) Navigate to **Manage > Software Management**.
- b) Click **Add Image**.
- c) In the **Add Software Image** window, select whether the image is **Remote** on a web server or **Local** on your machine.

In both cases, the image will be a file ending with `.iso`.

- **Remote:** Provide the **URL** to the image you downloaded in the first step.
- **Local:** Click **Choose file** and navigate to the local folder where you downloaded the image.

- d) Click **Add** to add the image.

Nexus Dashboard then downloads the upgrade image and starts processing the image, and goes through a number of preparation and validation stages to ensure successful upgrade. This may take several minutes to complete.

**Note**

See [Troubleshooting upgrades, on page 237](#) for more information on the validation checks that occur during this point of the upgrade and how to deal with upgrade issues that might arise.

- e) After the validation is complete, then the **Install** button appears in the card in the **Software Management** page. Click **Install** to install the software and go through the upgrade process.

The installation progress window is displayed. You can navigate away from this screen while the update is in progress.

This step may take up to 60 minutes or more, depending on the number of nodes in the cluster, during which the nodes will reboot and the GUI will not be accessible. Nexus Dashboard goes through several stages:

- Install Release Firmware
- Disable Services
- Shutdown Infrastructure services
- Update Platform Services
- Enable Infrastructure Services
- Enable Services

You can click on the **Details** link to see the progress and the various stages of the upgrade.

**Note**

If you see any issues during the upgrade process, such as a possible indexing issue, refer to [Prerequisites and guidelines for upgrading an existing Nexus Dashboard cluster, on page 221](#) for more information and possible workarounds.

After the process above is complete, you should be upgraded to Nexus Dashboard 4.3.1, where Nexus Dashboard and the individual services are packaged together as a single, unified product.

**Note**

Depending on the cluster format and the number of cluster nodes that you have deployed, certain features (such as controller, orchestrator, or telemetry) might not be available. Review the information in the [Nexus Dashboard Capacity Planning tool](#) to verify what features would be available for your cluster installation.

**Step 5** After the node upgrade tasks are completed, verify that the nodes are healthy and you can log into the UI.

Once the upgrade process completes, you can view the Nexus Dashboard UI as you typically would.

You can check the **Overview** page for overall system health and the **Admin > System Software** page to see the current Running version.

**Step 6** Review the cluster name changes in DNS.

Beginning with Nexus Dashboard release 4.3.1, for fresh installs and upgrades, the cluster DNS will be hardcoded to `nd-cluster.case.local`. When upgrading from previous releases to Nexus Dashboard release 4.3.1, the DNS domain name that you see based on the cluster name will change after the upgrade, and you will see the fixed cluster name (`nd-cluster.case.local`) rather than the cluster name within that DNS domain.

### What to do next

Go to [Post-upgrade information and tasks, on page 233](#) to perform the necessary post-upgrade tasks.

## Post-upgrade information and tasks

This section provides information about changes and tasks that you must complete after upgrading from Nexus Dashboard release 3.2.2 to Nexus Dashboard 4.3.1.

- [Perform multi-cluster connectivity post-upgrade tasks, on page 233](#)
- [Re-upload switch firmware images, on page 234](#)
- [Address Anomalies issues, on page 234](#)
- [Set device credentials, on page 234](#)
- [Migrate older cluster types to new 3 node virtual cluster \(data\) cluster type, on page 234](#)
- [Redeploy the telemetry configuration, on page 235](#)
- [SNMP-related tasks, on page 236](#)
- [Re-register ACI fabrics, on page 237](#)
- [View historical Performance Manager \(PM\) data, on page 237](#)
- [View system update notifications, on page 237](#)

### Perform multi-cluster connectivity post-upgrade tasks

If you had any of these configurations in pre-4.3.1. Nexus Dashboard:

- If you had Nexus Dashboard Insights onboarding remotely-created NDFC fabrics
- If you were using One Manage to manage and monitor multiple NDFC or NDI clusters
- If you had multi-cluster fabric groups with multiple NDFC clusters

If these clusters were already connected in the previous release using multi-cluster connectivity, then, on the primary cluster, you will have to re-register all the clusters that you upgraded to Nexus Dashboard release 4.3.1.

In addition, if you had NDI and NDO integrations in the previous release, this is not supported in Nexus Dashboard release 4.3.1, and in order to leverage the NDO integration, you will have to federate NDI and NDO clusters in Nexus Dashboard release 4.3.1.

If you had multi-cluster connectivity configured in your Nexus Dashboard release 3.2.2 system, after you have completed the upgrade, you will need to re-enable multi-cluster connectivity. This task is applicable only when you have NX-OS fabrics in a co-location environment. See [Connecting Clusters](#) for more information.

### Re-upload switch firmware images

Switch firmware images that were uploaded to Nexus Dashboard in release 3.2.2 will not be carried over when you upgrade to Nexus Dashboard 4.3.1. After upgrading to Nexus Dashboard 4.3.1, re-upload those switch firmware images:

1. Navigate to **Manage > Fabric Software > NX-OS/IOS-XE > Images**.
2. From the **Actions** drop-down list, select **Upload** and re-upload the necessary switch firmware images.

See [Managing Your Fabric Software](#) for more information.

### Address Anomalies issues

After you have completed the upgrade, check the **Anomalies** area for any issues and check for requests to complete the **Recalculate and Deploy** in some NX-OS fabrics..

### Set device credentials

If you had co-hosted NX-OS fabrics configured in the Nexus Dashboard release 3.2.2 system, follow these procedures to make sure the appropriate device credentials are configured after you've upgraded to Nexus Dashboard release 4.3.1.

1. In your upgraded Nexus Dashboard 4.3.1 system, navigate to **Manage > Device Credentials**.
2. Review the information displayed in the **Device Credentials** area.  
You should see **Not Set** displayed in red text in the **Device Credentials** area.
3. In the **Device Credentials** area, click **Set**.
4. In the **Set Default Credentials** page, enter the necessary information.
  - Username, Password, and Confirm password: Enter the necessary username and password information.
  - Robot: Choose the **Robot** checkbox to set the robot credentials, if necessary.

Then click **Save**.

You are returned to the **Set Default Credentials** page, and the text **Default Set** is now displayed in blue.

### Migrate older cluster types to new 3 node virtual cluster (data) cluster type

- These pre-4.3.1 cluster types are not supported as greenfield deployments in Nexus Dashboard release 4.3.1 but are supported when upgrading to release 4.3.1:
  - 3-node virtual cluster (app node with 1.5TB storage)
  - 5-node virtual cluster (app node with 1.5TB storage)
- The 5-node virtual cluster (app 500G) cluster type is not supported as a greenfield deployment in Nexus Dashboard release 4.3.1 but is supported when upgrading to release 4.3.1. However, this cluster type will not be supported as an upgrade option in the next release, so you will have to migrate to the new cluster type `3 node virtual cluster (data)` or the 3-node physical cluster using the procedures below.

If you want to migrate from those older cluster types to the new cluster type `3 node virtual cluster (data)` or to the 3-node physical cluster, you can migrate to the new cluster type using these procedures:

1. Upgrade to Nexus Dashboard release 4.3.1, keeping the older cluster type as-is, using the procedures provided in this chapter.
2. After you have upgraded to Nexus Dashboard release 4.3.1, perform a backup of your older cluster type. See [Backing Up and Restoring Your Nexus Dashboard](#) for more information.
  - You can use either a **Config-only** or **Full** backup option when backing up the older cluster type.
  - Verify that you have successfully retrieved the backup from the older cluster before proceeding.
3. Shut down the cluster with the older cluster type.
4. Perform a greenfield deployment of the new cluster type `3 node virtual cluster (data)`.
  - Reuse the name of the older cluster for the new cluster.
  - You can deploy a virtual data cluster or a physical cluster.
  - You can reuse the IP addresses from the older cluster type, or you can use new IP addresses in the new cluster and restore with a check in the **Ignore External Service IP Configuration** check box, as described in [Backing Up and Restoring Your Nexus Dashboard](#).
5. Restore the backup of the older `3 node virtual cluster (app)` **OR** `5 node virtual cluster (app)` to the newer `3 node virtual cluster (data)` or to the 3-node physical cluster. See [Backing Up and Restoring Your Nexus Dashboard](#) for more information.

### Redeploy the telemetry configuration

After upgrading to Nexus Dashboard release 4.3.1, the persistent IP addresses for handling software telemetry streaming from the NX-OS fabrics would have changed. To resume telemetry operations after the upgrade to Nexus Dashboard 4.3.1, you must redeploy the telemetry configuration.

1. Navigate to the **System Status** page.  
**Admin > System Status**
2. Click **Telemetry**, then click the **Fabrics** tab in the **Telemetry status** area.
3. Review the information displayed in the **Fabrics** table.

For NX-OS fabrics, you will notice that one or more fabrics listed in the **Fabrics** table will show a status of **Pending updates** in the **Telemetry config status** column, even though those same fabrics had a status of **OK** prior to the upgrade.

In this state, telemetry streaming will not be functional, and the status of individual features (such as software telemetry, flow collection, and switch status) are not valid and should be ignored.



**Note** Some fabrics in the **Fabrics** table might show a status of **OK** in the **Telemetry config status** column, while other fabrics might show a status of **Pending updates**. If you were to click the **Switches** tab in the **Telemetry status** area, you might see switches with incorrect green **OK** or **Success** status entries in the telemetry columns, even though those switches are associated with the fabrics that show a status of **Pending updates** in the **Fabrics** table. This is incorrect status information displayed at the switch level for those fabrics and should be ignored.

4. With the **Fabrics** tab selected in the **Telemetry status** area, click **Redeploy telemetry**.

Click **Confirm** in the confirmation page.

After you click **Confirm**, the status of the individual features will change to **Enable in Progress**, and the cumulative **Telemetry Config Status** will update to **In Progress**.

5. After you click **Confirm**, you are returned to the **Fabrics** tab in the **Telemetry** page under **System Status**. Click **Refresh** in the upper right corner of the page.

The telemetry status might be shown incorrectly immediately after you click **Confirm** in the confirmation page, but will show the correct telemetry status after you click **Refresh**.

6. Review the information displayed in the **Fabrics** table again.

After clicking **Refresh**, you should see that the status change in these areas:

- The status of the individual features will change to **Enable in Progress**, and
- The status for the fabric in the **Telemetry config status** column will change from **Pending updates** to **In Progress**. After several minutes, the status for the fabric will change to **OK** in the **Telemetry config status** column, either automatically or after you click **Refresh** again.

After the operation completes, the individual feature statuses will be:

- `Enabled` if the configuration push is successful for all switches,
- `Enable Fail` if it fails for any switch, or
- `Enable Pending` if change control mode is enabled. In that case, apply the change control ticket explicitly through **Manage > Change control**. See [Using Change Control and Rollback in Your Nexus Dashboard](#) for more information.

The cumulative **Telemetry Configuration Status** will then display as:

- `OK` if the configuration is successful for all switches,
- `Not OK` if it fails or is pending in change control mode for all switches, or
- `Partial OK` if it succeeds for some switches and fails or is pending in change control mode for others.

### SNMP-related tasks

- **Review changes in SNMP server users:** In releases prior to Nexus Dashboard release 4.3.1, SNMP server users that were configured on managed NX-OS switches without passwords as part of the intent, such as the example shown below:

```
snmp-server user Demo_CMDv5 vdc-operator
snmp-server user DemoOps_admin vdc-operator
snmp-server user DemoOps_admin network-admin
```

were not shown in the diffs during **Recalculate and Deploy** operations. These often went unnoticed, especially in environments where the switches relied on remote authentication methods such as TACACS+.

Once you upgrade to Nexus Dashboard release 4.3.1, these differences are now correctly detected and displayed in the GUI as expected diffs. After the upgrade to release 4.3.1, you must now either:

- Push these configurations to bring the switches into sync, or
- Remove these SNMP user entries from the intent if they are no longer needed.

- **Correct any invalid SNMP values:** Nexus Dashboard validates the SNMP listener IP and recipient email entries in release 4.3.1. If you had invalid values in those two fields in NDFC before the upgrade to Nexus Dashboard release 4.3.1, then you might see an `Input is not valid` error in the **SNMP listener IPs and ports for forwarded alarms data** field in this page:

**Admin > System Settings > Fabric management > Advanced settings > Alarms**

Fix the invalid entries in this field and click **Save** to address this issue.

### Re-register ACI fabrics

If you have onboarded ACI fabrics that are running on a release earlier than release 6.1.4, after you upgrade the cluster from 3.x to 4.x, you will have to re-register the onboarded ACI fabrics to use the ACI-to-Nexus Dashboard cross-launch functionality, as described in [Re-register clusters](#).

### View historical Performance Manager (PM) data

After you upgrade Nexus Dashboard from release 3.2.2 to release 4.3.1, if you enable Telemetry on the fabrics, the historical Performance Manager (PM) data will not be displayed. Instead, the system will start collecting PM data fresh from that point forward.

To view the historical PM data, disable Telemetry on the affected fabrics. The older PM data will now re-appear.

### View system update notifications

Beginning with release 4.2.1, Nexus Dashboard provides automated alerts to notify you when a new recommended software version is available. See the "System update notifications" section in [Exploring Your Nexus Dashboard](#) for more information.

## Troubleshooting upgrades

After all the nodes restart during new image activation stage described in the previous section, you may log in to the GUI to check the status of the upgrade workflows. Initially, you can see the bootstrap process similar to the initial cluster deployment and once the nodes come up, you can see additional information about service activation in the GUI's **Overview** page.

In case the upgrade fails for any reason, the GUI will display the error and additional workaround steps. For example, you might then see an error message, along with the remedy, similar to this:

```
Failed to activate
```

```
Upgrade failed while shutting down the cluster: Operation Timedout, last status: Operation Timedout
```

```
Please login to one of the primary nodes as 'rescue-user' and follow the steps provided by the upgrade recovery helper by invoking following command: 'acs upgrade recover Cluster Shutdown'. If the issue persists, please contact Cisco TAC for assistance.
```

If an issue persists, click **Admin** to access Tech support. See [Working with Cisco Tech Support](#) for more information.





## CHAPTER 11

# Upgrading a Nexus Dashboard 4.1.1 Cluster to This Release



**Note** The procedures in this chapter are applicable if you are upgrading from Nexus Dashboard release **4.1.1** to Nexus Dashboard release 4.3.1.

If you are upgrading from Nexus Dashboard release 3.2.2 to Nexus Dashboard release 4.3.1, follow the upgrade procedures provided in [Upgrading a Nexus Dashboard 3.2.2 Cluster to This Release, on page 221](#).

- [Prerequisites and guidelines for upgrading an existing Nexus Dashboard cluster, on page 239](#)
- [Nexus Dashboard pre-upgrade validation script, on page 241](#)
- [Supported upgrade paths, on page 243](#)
- [Upgrade Nexus Dashboard, on page 245](#)
- [Troubleshooting upgrades, on page 247](#)

## Prerequisites and guidelines for upgrading an existing Nexus Dashboard cluster

Before you upgrade your existing Nexus Dashboard cluster:

- Ensure that you have read the target release's [Release Notes](#) for any changes in behavior, guidelines, and issues that may affect your upgrade.
- Before upgrading to Nexus Dashboard release 4.3.1:
  - Make sure your NTP and DNS services are configured. At least one NTP and DNS are required for the system to upgrade successfully.
  - For LAN deployments, verify that the management network and data network are in different subnets. The upgrade will fail if the management network and data network are not in different subnets.

Note that this restriction does not apply for SAN deployments (management network and data network can be in same subnet for SAN deployments).

- We highly recommend that you use the Nexus Dashboard Preupgrade Validation script before performing any Nexus Dashboard upgrades. See [Nexus Dashboard pre-upgrade validation script, on page 226](#) for more information.
- Verify that the `acs health` is healthy.
  1. Access the Nexus Dashboard using `ssh -l rescue-user {management-ip-of-nd}`.
  2. Issue the `acs health` command.

The output from the `acs health` command should show that all components are healthy:

```
rescue-user@node1:~$ acs health
=====
Status
=====
All components are healthy
```

- Ensure that your current Nexus Dashboard cluster is healthy.
 

You can check the system status on the **Overview** page of the Nexus Dashboard's **Admin Console** or by logging in to one of the nodes as `rescue-user` and ensuring that the `acs health` command returns `All components are healthy`.
- If you have a multi-cluster federated Nexus Dashboard deployment, you must upgrade the primary cluster of the multi-cluster group to Nexus Dashboard version 4.3.1 or later before upgrading the local cluster.
- Review these important guidelines about how backup and restore processes affect your upgrade:
  - Ensure that you perform a backup of your Nexus Dashboard cluster before upgrading and you store the backup file in a safe place. To perform a backup, refer to [Backing Up and Restoring Your Nexus Dashboard](#).
  - An upgrade will not proceed if the most recent backup had a failure. Make sure you have a successful backup before proceeding with the upgrade. If you are unable to perform a successful backup and cannot upgrade, contact [Cisco Technical Assistance Center \(TAC\)](#) for support.
  - You will not be able to perform an upgrade if a restore has failed. After the restore has completed, click **View History** to navigate to the **History** area in the **Backup and Restore** page, as described in the "Restore Nexus Dashboard configurations" section in the [Unified Backup and Restore for Nexus Dashboard and Services](#).

The page should display **Success** in the **Status** column for the restore process. If you see any value other than **Success** in the **Status** column, redeploy the cluster and get a **Success** value for the restore process before attempting the upgrade again.

- If you are upgrading a physical Nexus Dashboard cluster, ensure that the nodes have the minimum supported CIMC version for the target Nexus Dashboard release.

Supported CIMC versions are listed in the [Nexus Dashboard Release Notes](#) for the target release.

The CIMC upgrade is described in detail in the "Troubleshooting" article in the Nexus Dashboard [documentation library](#).

- If you are upgrading a virtual Nexus Dashboard cluster, Nexus Dashboard will enforce these checks:
  - A check of the HDD latency to verify that it is <30ms. If the HDD has a higher latency, the upgrade will fail.

- A check of the network latency to verify that it is <50ms within cluster nodes. If the network has a higher latency, the upgrade will fail.
- If you are upgrading a virtual Nexus Dashboard cluster deployed in VMware ESX, ensure that the ESX version is still supported by the target release.

This release supports VMware ESXi 7.0, 7.0.1, 7.0.2, 7.0.3, 8.0, 8.0.2, 8.0.3.

**Note**

If you need to upgrade the ESX server, you must do that before upgrading your Nexus Dashboard. ESX upgrades are outside the scope of this document, but in short:

1. Upgrade one of the ESX hosts as you typically would with your existing Nexus Dashboard node VM running.
2. After the host is upgraded, ensure that the Nexus Dashboard cluster is still operational and healthy.
3. Repeat the upgrade on the other ESX hosts one at a time.
4. After all ESX hosts are upgraded and the existing Nexus Dashboard cluster is healthy, proceed with upgrading your Nexus Dashboard to the target release as described in this document.

- Nexus Dashboard does not support platform downgrades.

If you want to downgrade to an earlier release, you will need to deploy a new cluster.

**Updates for Nexus Dashboard release 4.3.1**

These are the updates for 4.3.1.

- OVA template changes for Nexus Dashboard release 4.3.1:
  - CPU and memory reservations are set by default.
  - Removed customizing disk sizes; you select App or Data.
  - Default deployment flavor changed from App to Data.
- Upgrade history is now added in Nexus Dashboard release 4.3.1.
- Retry buttons added for upgrade failures and retries.

## Nexus Dashboard pre-upgrade validation script

We highly recommend that you use the Nexus Dashboard pre-upgrade validation script before performing any Nexus Dashboard upgrades. The Nexus Dashboard pre-upgrade validation script performs various checks for known issues that have been identified to affect the success of a Nexus Dashboard upgrade. The script is continuously updated and maintained in an effort to mitigate any new upgrade-related issues that are detected in the field.

- Prior to Nexus Dashboard release 4.3.x, the Nexus Dashboard pre-upgrade validation script was a standalone Python script that you would [download](#) and run from an external host over SSH to perform health checks across nodes.
- Beginning with Nexus Dashboard release 4.3.x, the Nexus Dashboard pre-upgrade validation script is now a signed plugin process, where you download the Cisco-signed `.ndp` (Nexus Dashboard plugin) file, which Nexus Dashboard validates and runs locally with elevated privileges.

### Guidelines and limitations: Pre-upgrade validation script commands

These are the guidelines and limitations for the pre-upgrade validation script commands:

- Only one plugin can be installed at a time. If you want to download the Nexus Dashboard pre-upgrade validation script plugin file and you already have a plugin file installed, you must remove the installed plugin first, then you can download and run another Nexus Dashboard pre-upgrade validation script plugin file. See [Pre-upgrade validation script commands, on page 242](#) for more information.

### Pre-upgrade validation script commands

The pre-upgrade validation script commands are `acs` CLI commands, similar to the Nexus Dashboard `acs` CLI commands that are described in the [Cisco Nexus Dashboard Troubleshooting](#) article. Log into any of the nodes in the Nexus Dashboard cluster to run these pre-upgrade validation script commands `acs` CLI commands.

These are the commands that are available for the new Nexus Dashboard `.ndp` version of the pre-upgrade validation script:

- To download the Nexus Dashboard pre-upgrade validation script plugin file, enter:

```
acs plugin download <url>
```

where `<url>` is the plugin download URL. This command downloads, verifies, and installs the plugin file.

The supported methods are positional arguments, and the supported download formats are:

- **HTTP download:** `http://<server>/<path>`. For example:

```
acs plugin download http://192.168.1.100/my-plugin.ndp
```

- **HTTPS download:** `https://<server>/<path>`. For example:

```
acs plugin download https://server.example.com/plugins/my-plugin.ndp
```

- **SFTP from remote secure storage:** `sftp://<server>:/<path>`. For example:

```
acs plugin download sftp://sftp-server.example.com:/export/plugins/my-plugin.ndp
```

- **SCP from remote secure storage:** `scp://<server>:/<path>`. For example:

```
acs plugin download scp://scp-server.example.com:/export/plugins/my-plugin.ndp
```

**Note**

- In order to use the SCP and SFTP download formats, you must first configure the remote storage. Refer to [Create a remote storage location](#) for those procedures and choose the **SFTP/SCP Server** option in the **Remote Storage Location Type** field.
- SCP (`scp://`) also supports SFTP for backwards compatibility.

If you already have a plugin file installed, you will see the message `plugin already installed, run 'acs plugin delete' first`. Since you can only have one installed plugin file at a time, determine if you want to replace the existing plugin file:

- Enter `acs plugin list` to determine the version for the plugin that's currently installed. You can choose not to install the new plugin file if you see that the version is the same.
- If you want to replace the existing plugin file, remove the installed plugin file using the `acs plugin delete` command as described below, then attempt to download the new plugin file again.
- To run the Nexus Dashboard pre-upgrade validation script plugin file, enter:  

```
acs plugin run
```

If you do not have a plugin installed, the output returns a `no plugin installed` message.
- To show plugin file information, enter:  

```
acs plugin list
```

  - If you have a plugin installed, the output returns the name, version, and `signed_by` information.
  - If you do not have a plugin installed, the output returns a `no plugin installed` message.
- To remove an installed plugin file, enter:  

```
acs plugin delete
```

Since you can only have one installed plugin file at a time, use this command to remove an already installed plugin file if you are attempting to download a new plugin file.

## Supported upgrade paths

As described in [Nexus Dashboard deployment overview, on page 5](#), in earlier releases, Nexus Dashboard shipped with only the platform software and no services included, which you would then download, install, and enable separately after the initial platform deployment. In addition, Nexus Dashboard release 3.1.1 introduced a tighter coupling between the Nexus Dashboard and individual services with only a single version of each service compatible with each version of the platform. As a result, as long as you were on the minimum required version of the Nexus Dashboard software, you could upgrade both the platform and all currently enabled services directly to Nexus Dashboard release 3.1x and 3.2x.

Beginning with Nexus Dashboard release 4.1.1, the platform and the individual services have been unified into a single product, which means that you no longer deploy, configure, or upgrade the services separately.

The following table provides a few example scenarios for specific deployment combinations:

Table 29:

| Current Nexus Dashboard Release | Compatible Services<br>(depending on form factor and cluster size, you may have one or more of these services currently enabled) | Upgrade Workflow                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                           |
|---------------------------------|----------------------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 4.1.1                           | N/A                                                                                                                              | Upgrade directly to release 4.3.1 as described in the following section.                                                                                                                                                                                                                                                                                                                                                                                                                                                   |
| 3.2.2                           | Fabric Controller: 12.2(3)<br>Orchestrator: 4.4(2)<br>Insights: 6.5(2)                                                           | Upgrade directly to release 4.3.1 as described in <a href="#">Upgrading a Nexus Dashboard 3.2.2 Cluster to This Release, on page 221</a> .<br><br>All services are unified under a single Nexus Dashboard product in release 4.3.1.                                                                                                                                                                                                                                                                                        |
| 3.2.1                           | Fabric Controller: 12.2(2)<br>Orchestrator: 4.4(1)<br>Insights: 6.5(1)                                                           | <ol style="list-style-type: none"> <li>1. Upgrade the Nexus Dashboard platform to release 3.2.2 as described in <a href="#">Nexus Dashboard Deployment Guide, Release 3.2.x</a>.<br/>All services will be automatically upgraded along with the platform.</li> <li>2. Upgrade from release 3.2.2 to release 4.3.1 as described in <a href="#">Upgrading a Nexus Dashboard 3.2.2 Cluster to This Release, on page 221</a>.<br/>All services are unified under a single Nexus Dashboard product in release 4.3.1.</li> </ol> |
| 3.1.1                           | Fabric Controller: 12.2(1)<br>Orchestrator: 4.3(x)<br>Insights: 6.4(1)                                                           | <ol style="list-style-type: none"> <li>1. Upgrade the Nexus Dashboard platform to release 3.2.2 as described in <a href="#">Nexus Dashboard Deployment Guide, Release 3.2.x</a>.<br/>All services will be automatically upgraded along with the platform.</li> <li>2. Upgrade from release 3.2.2 to release 4.3.1 as described in <a href="#">Upgrading a Nexus Dashboard 3.2.2 Cluster to This Release, on page 221</a>.<br/>All services are unified under a single Nexus Dashboard product in release 4.3.1.</li> </ol> |
| 3.0.1                           | Fabric Controller: 12.1(3)<br>Orchestrator: 4.2(x)<br>Insights: 6.3(1)                                                           | <ol style="list-style-type: none"> <li>1. Upgrade the Nexus Dashboard platform to release 3.2.2 as described in <a href="#">Nexus Dashboard Deployment Guide, Release 3.2.x</a>.<br/>All services will be automatically upgraded along with the platform.</li> <li>2. Upgrade from release 3.2.2 to release 4.3.1 as described in <a href="#">Upgrading a Nexus Dashboard 3.2.2 Cluster to This Release, on page 221</a>.<br/>All services are unified under a single Nexus Dashboard product in release 4.3.1.</li> </ol> |

| Current Nexus Dashboard Release | Compatible Services<br>(depending on form factor and cluster size, you may have one or more of these services currently enabled) | Upgrade Workflow                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                            |
|---------------------------------|----------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 2.3.2 and earlier               | Fabric Controller: 12.1(2) or earlier<br>Orchestrator: 4.1(x) or earlier<br>Insights: 6.2(x) or earlier                          | <ol style="list-style-type: none"> <li>1. Upgrade the Nexus Dashboard platform to release 3.1.1 as described in <a href="#">Nexus Dashboard Deployment Guide, Release 3.1.x</a><br/>All services will be automatically upgraded along with the platform.</li> <li>2. Upgrade from release 3.1.1 to release 3.2.2 as described in <a href="#">Nexus Dashboard Deployment Guide, Release 3.2.x</a>.<br/>All services will be automatically upgraded along with the platform.</li> <li>3. Upgrade from release 3.2.2 to release 4.3.1 as described in <a href="#">Upgrading a Nexus Dashboard 3.2.2 Cluster to This Release, on page 221</a>.<br/>All services are unified under a single Nexus Dashboard product in release 4.3.1.</li> </ol> |

## Upgrade Nexus Dashboard

This section describes how to upgrade an existing Nexus Dashboard 4.1.1 cluster to the Nexus Dashboard 4.3.1 release.

### Before you begin

Ensure that you have completed the prerequisites described in [Prerequisites and guidelines for upgrading an existing Nexus Dashboard cluster, on page 239](#)

### Procedure

#### Step 1

In your Nexus Dashboard release 4.1.1 system, download the Nexus Dashboard 4.3.1 image.

- a) Browse to the Software Download page.

<https://software.cisco.com/download/home/286327743/type/286328258>

- b) Choose the Nexus Dashboard 4.3.1 release to download.
- c) Download the Nexus Dashboard image for the 4.3.1 release.

#### Note

- The upgrade process is the same for all Nexus Dashboard form factors and uses the Nexus Dashboard ISO image (nd-dk9.<version>.iso). In other words, even if you used the virtual form factors (such as the ESX .ova) or a cloud provider's marketplace for initial cluster deployment, you must still use the .iso image for upgrades.

- If the image fails to download completely, verify the network connectivity between the Nexus Dashboard and the image server. Check the proxy configuration under **Admin > System Settings > General > Proxy configuration**.

d) (Optional) Host the image on a web server in your environment.

**Note**

We recommend hosting the image on a server in your environment. When you upload the image to your Nexus Dashboard cluster, you will have an option to provide a direct URL to the image, which can significantly speed up the process.

**Step 2** Log in to your current Nexus Dashboard's **Admin Console** as an `Administrator` user.

**Step 3** Delete any older, non-active upgrade images from your cluster.

If this is the first time you're upgrading your cluster, you can skip this step.

- a) Navigate to **Manage > Software Management**.
- b) Click the trash icon on an upgrade image's tile to delete any older, non-active upgrade **Images**.
- c) Repeat this step for all older, non-active upgrade images.

**Step 4** Upload the new image to the cluster.

- a) Navigate to **Manage > Software Management**.
- b) Click **Add Image**.
- c) In the **Add Software Image** window, select whether the image is **Remote** on a web server or **Local** on your machine.

In both cases, the image will be a file ending with `.iso`.

- **Remote:** Provide the **URL** to the image you downloaded in the first step.
- **Local:** Click **Choose file** and navigate to the local folder where you downloaded the image.

d) Click **Add** to add the image.

Nexus Dashboard then downloads the upgrade image and starts processing the image, and goes through a number of preparation and validation stages to ensure successful upgrade. This may take several minutes to complete.

**Note**

See [Troubleshooting upgrades, on page 247](#) for more information on the validation checks that occur during this point of the upgrade and how to deal with upgrade issues that might arise.

e) After the validation is complete, these buttons appear in the card in the **Nexus Dashboard Releases** area in the **System Software** page.

- **Factory install:** The **Factory install** option wipes all data, including bootstrap configurations. You will have to re-bootstrap the cluster afterward.
- **Upgrade:** Click the **Upgrade** option if you are upgrading from a previous release. You do not have to perform a **Factory install** before upgrading.

Click the appropriate button to install or upgrade the software.

The installation progress window is displayed. You can navigate away from this screen while the update is in progress.

This step may take up to 60 minutes or more, depending on the number of nodes in the cluster, during which the nodes will reboot and the GUI will not be accessible. Nexus Dashboard goes through several stages:

- Install Release Firmware
- Disable Services
- Shutdown Infrastructure services
- Update Platform Services
- Enable Infrastructure Services
- Enable Services

You can click on the **Details** link to see the progress and the various stages of the upgrade.

**Note**

If you see any issues during the upgrade process, such as a possible indexing issue, refer to [Prerequisites and guidelines for upgrading an existing Nexus Dashboard cluster, on page 239](#) for more information and possible workarounds.

After the process above is complete, you should be upgraded to Nexus Dashboard 4.3.1.

**Note**

Depending on the cluster format and the number of cluster nodes that you have deployed, certain features (such as controller, orchestrator, or telemetry) might not be available. Review the information in the [Nexus Dashboard Capacity Planning tool](#) to verify what features would be available for your cluster installation.

**Step 5** After the node upgrade tasks are completed, verify that the nodes are healthy and you can log into the UI.

Once the upgrade process completes, you can view the Nexus Dashboard UI as you typically would.

You can check the **Overview** page for overall system health and the **Admin > System Software** page to see the current Running version.

**Step 6** Review the cluster name changes in DNS.

Beginning with Nexus Dashboard release 4.3.1, for fresh installs and upgrades, the cluster DNS will be hardcoded to `nd-cluster.case.local`. When upgrading from previous releases to Nexus Dashboard release 4.3.1, the DNS domain name that you see based on the cluster name will change after the upgrade, and you will see the fixed cluster name (`nd-cluster.case.local`) rather than the cluster name within that DNS domain.

---

## Troubleshooting upgrades

After all the nodes restart during new image activation stage described in the previous section, you may log in to the GUI to check the status of the upgrade workflows. Initially, you can see the bootstrap process similar to the initial cluster deployment and once the nodes come up, you can see additional information about service activation in the GUI's **Overview** page.

In case the upgrade fails for any reason, the GUI will display the error and additional workaround steps. For example, you might then see an error message, along with the remedy, similar to this:

```
Failed to activate
```

```
Upgrade failed while shutting down the cluster: Operation Timedout, last status: Operation
Timedout
```

```
Please login to one of the primary nodes as 'rescue-user' and follow the steps provided by
```

the upgrade recovery helper by invoking following command: 'acs upgrade recover Cluster Shutdown'. If the issue persists, please contact Cisco TAC for assistance.

If an issue persists, click **Admin** to access Tech support. See [Working with Cisco Tech Support](#) for more information.



## CHAPTER 12

# Migrating From DCNM to ND

- [Migrating from DCNM to ND, on page 249](#)

## Migrating from DCNM to ND

You cannot upgrade directly from DCNM 11.5(4) to Nexus Dashboard 4.3.1. Use this workflow instead to upgrade from DCNM 11.5(4) to Nexus Dashboard 4.3.1:

1. First upgrade from DCNM 11.5(4) to Nexus Dashboard 4.1.1 using the procedures provided in the "Migrating From DCNM to ND" section in [Cisco Nexus Dashboard Deployment and Upgrade Guide, Release 4.1.x](#).
2. Then upgrade from Nexus Dashboard 4.1.1 to Nexus Dashboard 4.3.1 using the procedures provided in [Upgrading a Nexus Dashboard 4.1.1 Cluster to This Release, on page 239](#).



### Note

There is a known issue where the Nexus Dashboard Orchestrator (NDO) app is disabled after upgrading from DCNM 11.5 to Nexus Dashboard 4.1.1. You will have to manually re-enable the NDO (Orchestrator) app before upgrading from Nexus Dashboard 4.1.1 to Nexus Dashboard 4.3.1; otherwise, you will see an upgrade pre-check failure.

1. Enable the **Display advanced settings and options for TAC support** feature (**Admin > System Settings > General > Advanced Settings > Display advanced settings and options for TAC support**).

You have to enable the **Display advanced settings and options for TAC support** feature in order to see the **Features** tab in the next step.

2. Manually re-enable the NDO (Orchestrator) app (**Admin > System Status > Features**).

You can then disable the **Display advanced settings and options for TAC support** feature again, if necessary.





## CHAPTER 13

# Migrate from NDB to ND

- [Migrate from NDB to ND, on page 251](#)

## Migrate from NDB to ND

If you are currently using the NDB centralized version 3.10.4 (or 3.10.5.x) and intend to upgrade, you can upgrade to the NDB version which is available as part of ND 4.2.

Use this procedure for migrating from the centralized deployment of Nexus Data Broker (NDB) to Nexus Dashboard (ND).

The supported paths for migrating to ND 4.2 are:

- Migration from NDB 3.10.4 or 3.10.5.x to ND 4.2.
- For NDB versions prior to release 3.10.4, migration to ND 4.2 is a two-step process. NDB versions prior to 3.10.4 must first be upgraded to 3.10.4 (or 3.10.5.x). The supported NDB versions for this upgrade are: 3.8.x, 3.9.x, 3.10.1, 3.10.2, 3.10.3.

Refer to the respective [Nexus Data Broker Deployment Guide](#) for the detailed procedure for upgrading to NDB release 3.10.4 or 3.10.5.x.

Example: For upgrading NDB 3.8.1 to ND 4.2, first upgrade to NDB 3.10.4 (or 3.10.5.x), and then migrate to ND 4.2.

### Before you begin

1. Log in to the Nexus Data Broker GUI, and download the configuration file using the standard procedure.  
On the NDB GUI, navigate to **Administration** > **Backup Restore** > **Actions**. You can back up the configuration to your local machine or to a server.
2. Convert the NDB configuration file (which you have downloaded using step 1) to an ND-compatible file using the `backup_converter.sh` file located on [Github](#).
3. (optional) Create an encryption key using `./backup_converter.sh -f {zip_file} -v {ND_version} -k {password}`. See *Run the Script* step ([Github](#)). This encryption key cannot be changed later. Ensure to save the encryption key as you will need to enter the same key while restoring the backup.

4. Follow the steps detailed in the `steps_to_generate_ND_format_backup.txt` file (available on [Github](#)). The file includes a detailed procedure for converting the NDB configuration file to an ND-compatible file. The ND compatible file is named as `cisco-nddb-backup.tar.gz`.

## Procedure

---

- Step 1** Log in to the Nexus Dashboard GUI.
- Step 2** Navigate to **Admin > Backup and Restore**.
- Step 3** On the **Backup and Restore** screen, click the **Restore** button.
- Step 4** On the **Restore** screen, enter:
- a) **Source**: the location of the backup file.
  - b) **Backup file**: the converted ND compatible file.
  - c) **Encryption key**: enter the custom encryption key which you had earlier created. If you have not set an encryption key, the default value, `cisco123`, is assigned.
  - d) Click **Validate and upload**. The **Filename** field is automatically populated with the ND compatible file (`cisco-nddb-backup.tar.gz`).
  - e) Click **Next** after the validation is complete.
- The progress of the validation is indicated by the green bar, wait till 100% is displayed.
- Step 5** (on the next page of the Restore screen) Check the **Do not restore persistent IP configs from this backup** check box to override the persistent IP configuration.
- Ensure to check the check box, else the restore process will fail.
- Step 6** Click **Restore**.
- Step 7** Click **Restore** (again) on the **Begin Restore** confirmation screen.
- After the confirmation, the migration process is initiated. You can monitor the progress; the various stages are indicated by the progress bar.
- Step 8** Click **Close**.
- 

## What to do next

After the migration process, perform these steps:

1. On the ND GUI, navigate to **Manage > Fabrics**.
2. Select the required NDB fabric (NDB\_default or if you have created a slice, the *slice name* is the *fabric name*).
3. Click **Actions** (on the top right) > **Recalculate and deploy**.  
This updates the configuration on the NDB devices.
4. Check the device mode (**Manage > Inventory**). The mode is now displayed as *Normal* (during the migration process, the device mode is indicated as *Migration*).