



NVIDIA DGX SuperPOD: Next Generation Scalable Infrastructure for AI Factories

Reference Architecture

Featuring NVIDIA DGX Vera Rubin NVL72 Systems

Abstract

The NVIDIA DGX SuperPOD architecture has been designed to power next-generation AI factories with unparalleled performance, scalability, and innovation that supports all customers in the enterprise, finance, higher education, research, and the public sector. It is a physical twin of NVIDIA's own system used by research and development teams, meaning the DGX SuperPOD customer's infrastructure software, applications, and support have already been tested and vetted on the same architecture through NVIDIA's internal usage



This DGX SuperPOD Reference Architecture (RA) is based on DGX Vera Rubin NVL72 systems powered by Vera CPUs and Rubin GPUs. The RA discusses the components that define the scalable and modular architecture of DGX SuperPOD. DGX SuperPOD is built on the concept of scalable units (SU); each SU contains 8 DGX Vera Rubin NVL72 systems, which enables the rapid deployment of DGX SuperPOD of any size. The RA also includes details regarding the SU design and specifics of InfiniBand, NVLink network, Ethernet fabric topologies, storage system specifications, recommended rack layouts, and wiring guidelines.

This RA combines the latest NVIDIA technologies to help companies and industries develop their own AI factories. To achieve the most scalability, DGX

SuperPOD is powered by several key NVIDIA technologies and solutions, including:

- NVIDIA DGX Vera Rubin NVL72 systems: Powerful computational building block for AI and HPC.
- NVIDIA QUANTUM-X800 (XDR/800 Gbps) InfiniBand: High performance, low latency, and scalable network interconnect.
- NVIDIA Spectrum-4 (800 Gbps) Ethernet: High performance, low latency, and scalable ethernet connectivity for storage.
- NVIDIA [sixth-generation NVLink](#)® technology: a high-speed interconnect for CPU and GPU processors in accelerated compute systems to provide unprecedented performance for most demanding communication patterns.
- [NVIDIA Mission Control](#): a unified operations and orchestration software stack for managing AI factories.

DGX SuperPOD requires a supported storage solution from one of the validated and certified storage vendor from the [DGX Storage Program](#). A full list of supported storage vendors are listed on the [DGX SuperPOD landing page](#).

Contents

DGX SuperPOD Architecture	4
Key Components of DGX SuperPOD	6
<i>NVIDIA DGX Rack System.....</i>	<i>7</i>
DGX Vera Rubin NVL72 Compute Tray	7
NVLink Switch Tray	8
DGX Power Shelves	9
<i>NVIDIA InfiniBand Technology.....</i>	<i>9</i>
<i>NVIDIA Mission Control</i>	<i>10</i>
<i>Components.....</i>	<i>11</i>
<i>Design Requirements.....</i>	<i>12</i>
System Design.....	13
Compute Fabric.....	13
Ethernet Fabric	13
Out-of-Band Management Network.....	14
Storage Requirements	14
Network Fabrics.....	16
Multi-node NVLink Fabric (NVL6)	17
Compute Fabric.....	18
Ethernet Fabric	Error! Bookmark not defined.
Network Segmentation of the Ethernet Fabric.....	Error! Bookmark not defined.
High Performance Storage Architecture	26
Management Servers.....	29
DGX SuperPOD Software	30
Summary	31

DGX SuperPOD Architecture

The DGX SuperPOD architecture is a combination of DGX systems, InfiniBand and Ethernet networking, management nodes, and storage. Figure 1 shows the rack layout of a single SU (Scalable Unit). Each SU requires a Thermal Design Power (TDP) of ~2 Megawatts (MW). Generally, the data center should meet or exceed Uptime Institute Tier 3 design standards, or alternatively the TIA942-B Rated 3 or EN50600 Availability Class 3 design standards, including concurrent maintainability and no single point of failure.

Each DGX Vera Rubin NVL72 leverages liquid cooling to manage the amount of heat produced by each rack. Given DGX SuperPOD with DGX Vera Rubin NVL72 is a turnkey AI data center scale product, NVIDIA has provided [datacenter design guidance](#) and fully integrated operational technology (OT) integration with NVIDIA Mission Control software stack.



Figure 1. Sample Design for DGX SuperPOD Vera Rubin NVL72

This reference architecture showcases the scale-out capabilities of the DGX SuperPOD up to 128 rack configurations with 9216 GPUs.

Contact your NVIDIA representative for information regarding DGX SuperPOD solutions of four SUs or more.

Key Components of DGX SuperPOD

The DGX SuperPOD architecture is designed to meet the high demands of AI factories for the era of reasoning & reinforced learning. AI factories require specialized components like high-performance GPUs and CPUs, advanced networking, and cooling systems to support the intensive computational needs of AI workloads. These factories excel at AI reasoning—enabling faster, more accurate decision-making across industries. And using NVIDIA's end-to-end accelerated computing platform, they're optimized for energy efficiency while accelerating AI inference performance, helping enterprises deploy secure, future-ready AI infrastructure.

DGX SuperPOD is the integration of key NVIDIA components, as well as storage solutions from partners certified to work in a DGX SuperPOD environment. By leveraging the concept of Scalable Units (SUs, as defined below in this document), DGX SuperPOD reduces AI factory deployment times from months to weeks, which in turn reduces time-to-solution and time-to-market of next generation models and applications.

DGX SuperPOD is to be deployed on-premises, meaning the customer owns and manages the hardware. This can be within a customer's data center or co-located at a commercial data center, but the customer owns the hardware, the service it provides, and is responsible for their cluster infrastructure.

The key components of DGX Vera Rubin NVL72 SuperPOD are shown below, each component will be discussed in detail:

- NVIDIA DGX Vera Rubin NVL72 rack system
- NVIDIA QUANTUM-X800 InfiniBand
- NVIDIA Spectrum Ethernet
- NVIDIA Mission Control software

NVIDIA DGX Rack System

NVIDIA DGX Vera Rubin NVL72 is an AI powerhouse that enables enterprises to expand the frontiers of business innovation and optimization. Each DGX Vera Rubin NVL72 delivers breakthrough AI performance in a rack-scale, 72 GPU configuration. The NVIDIA Rubin GPU architecture provides the latest technologies that brings months of computational effort down to days and hours, on some of the largest AI/ML workloads.

To accommodate the high compute density per rack and at the same time to more efficiently use datacenter space, the DGX Vera Rubin NVL72 rack system employs a sophisticated cooling solution to manage the substantial heat generated by the powerful GPUs and other components. NVIDIA's Infrastructure Specialist team can provide end-to-end assistance for datacenter planning, system deployment, and bring-up services.

DGX Vera Rubin NVL72 Compute Tray

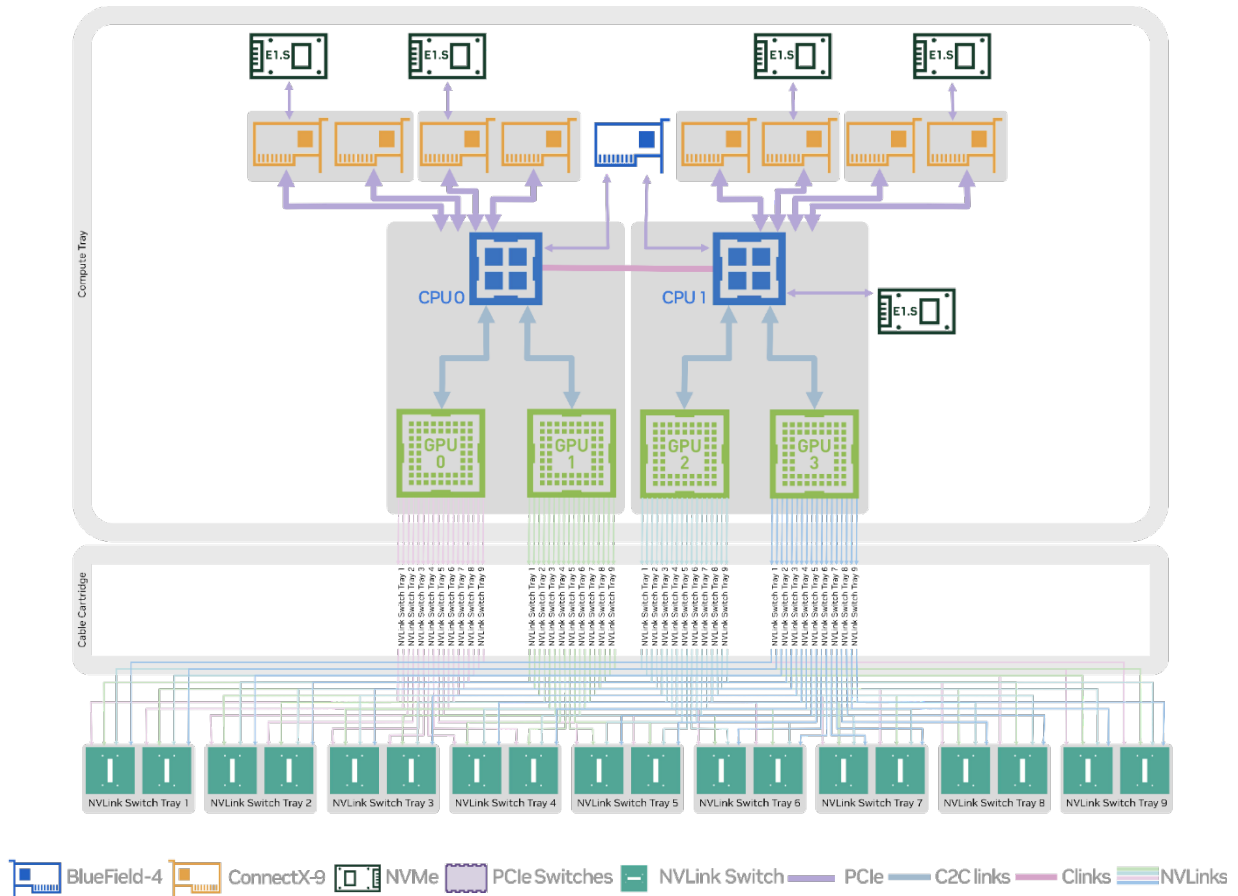
The compute nodes for DGX Vera Rubin NVL72 rack system are using the NVL72 topology, which contains 72 GPUs in a single NVLink domain. There are 18 compute nodes (aka. compute trays) in a DGX Vera Rubin NVL72 rack system. Each compute tray contains two Vera Rubin NVL72 Superchips, and each Superchip has two Rubin GPUs and one Vera CPU. A coherent chip-to-chip interconnect link, called NVLink-C2C, bridges the two Superchips and enables them to act as a single logical unit for a single OS instance to operate.

The compute tray integrates eight ConnectX-9 NICs, with 2x800 Gbps InfiniBand X-800 connectivity for the inter-racks Compute network, and one BlueField-4 NIC to support 2x400Gbps connectivity for the In-band Management and Storage networks. All network ports are located at the front-side of the rack, intentionally for the cold (or hot) aisle.

Each compute tray also provides a total of 4x 7.68TB E1.S NVMe as RAID0 for local storage and 1x 2TB E1.S NVMe for the OS image.

Figure 2 shows the block diagram of the Vera Rubin NVL72 compute tray with two Vera Rubin NVL72 Superchips, each chip combines two NVIDIA Rubin Tensor Core GPUs and one NVIDIA Vera CPU, connected over a 1,800GB/s ultra-low-power NVLink-C2C interconnect.

Figure 2. Vera Rubin NVL72 Compute Tray Block Diagram



NVLink Switch Tray

Each DGX Vera Rubin NVL72 system rack also features nine NVLink switches (aka. NVLink switch tray). All NVLink interconnections are done through the rear-side, blind-mate connectors to the cable cartridges within the rack.

DGX Power Shelves

The power shelf used for DGX Vera Rubin NVL72 SuperPOD has six 18.3kW PSUs and can deliver up to 110kW of power. There are four total power shelves in a single DGX Vera Rubin NVL72 rack system, with only three required for full rack power, allowing for N+N redundancy in the data center on every rack. The DGX Vera Rubin NVL72 power shelf is equipped with enhanced energy storage and additional bulk capacitors to support power smoothing and to reduce the peak load demands on the datacenter power infrastructure.

At the rear of the power shelf is a set of RJ45 ports used for power brakes and other out-of-band management capabilities, that are interconnected with other power shelves in the same rack. At the front of the power shelf is the BMC port. NVIDIA DGX Power shelves include deep integration into NVIDIA Mission Control Software for enhanced power monitoring and management.

NVIDIA InfiniBand Technology

InfiniBand is a high-performance, low latency, RDMA capable networking technology, proven over 20 years in the harshest compute environments to provide the best inter-node network performance. InfiniBand continues to evolve and lead data center network performance.

The latest generation InfiniBand, Quantum-X800, has a peak speed of 800 Gbps per direction with an extremely low port-to-port latency. It is backwards compatible with the previous generations of InfiniBand specifications. InfiniBand is more than just peak bandwidth and low latency. It provides additional features to optimize performance including adaptive routing (AR), collective communication with SHARP™, dynamic network healing with SHIELD™, and supports several network topologies including fat-tree, Dragonfly, and multi-dimensional Torus to build the largest network fabrics and computer systems possible.

NVIDIA Mission Control

The DGX Vera Rubin NVL72 SuperPOD Reference Architecture represents the best practices for building high-performance AI factories. There is flexibility in how these systems can be presented to customers and users. NVIDIA Mission Control software is used to manage all DGX Vera Rubin NVL72 SuperPOD deployments.

NVIDIA Mission Control is a sophisticated full-stack software solution. As an essential part of the DGX SuperPOD experience, it optimizes developer workload performance and resiliency, ensures continuous uptime with automated failure handling, and provides unified cluster-scale telemetry and manageability. Key features include full-stack resiliency, predictive maintenance, unified error reporting, data center optimizations, cluster health checks, and automated node management.

NVIDIA Mission Control software incorporates the same technology NVIDIA uses to manage thousands of systems in its own AI factories and provides a proven path to operating large-scale, production-grade AI infrastructure for organizations that need the highest performance and reliability.

DGX SuperPOD is to be deployed on-premises, meaning the customer owns and manages the hardware. This can be within a customer's data center or co-located at a commercial data center, but the customer owns the hardware, the service it provides, and is responsible for their cluster infrastructure as well as providing the building management system for integration.

Components

The hardware components of DGX Vera Rubin NVL72 SuperPOD are described in and the software components are shown in Table 2.

Table 1. DGX SuperPOD / 1 Scalable Unit infrastructure

Component	Technology	Description
8x Compute Racks	NVIDIA DGX Vera Rubin NVL72	The world's premier rack-scale, purpose-built AI systems featuring NVIDIA Vera-Rubin NVL72 modules. Vera Rubin NVL72 Superchips. NVL72 scale-out GPU interconnect and integrated NVLink switch trays.
NVLink Fabric	NVIDIA NVLink 6	NVLink Switches supports fast, direct memory access between GPUs on the same compute rack.
Compute Fabric	NVIDIA Q3400 InfiniBand Switch	Rail-optimized, non-blocking, full fat-tree network with eight Quantum X-800 connections per system for inter-rack GPU communications.
Storage and In-band Management Fabric	NVIDIA Spectrum 4 SN5600 Ethernet switches	The fabric is optimized to match peak performance of the configured storage array, built with 64 port 800 Gbps Ethernet switch providing high port density with high performance.
InfiniBand management	NVIDIA Unified Fabric Manager Appliance, Enterprise Edition	NVIDIA UFM combines enhanced, real-time network telemetry with AI powered cyber intelligence and analytics to manage scale-out InfiniBand data centers
NVLink Management and Ethernet Monitoring	NVIDIA NetQ	NVIDIA NetQ manages, operates the NVLink switches and provides real-time network telemetry to manage all NVLink infrastructure, as well as to the Ethernet switches and fabric.
Out-of-band (OOB) management network	NVIDIA SN2201 switch	48 x 1 Gbps Ethernet switch leveraging copper ports to minimize complexity

Table 2. DGX SuperPOD software components

Component	Description
-----------	-------------

Mission Control Software	NVIDIA Mission Control software delivers a full stack data center solution engineered for NVIDIA DGX SuperPOD with DGX Vera Rubin NVL72 infrastructure deployments. It integrates essential management and operational capabilities into a unified platform, thereby providing enterprise customers with simplified control over their NVIDIA DGX SuperPOD infrastructure deployments at scale. NVIDIA Mission Control also leverages NVIDIA Run:ai functionality, providing seamless workload orchestration.
NVIDIA AI Enterprise	NVIDIA AI Enterprise is an end-to-end, cloud-native software platform that accelerates data science pipelines and streamlines development and deployment of production-grade LLMs and other generative AI applications.
Magnum IO	Enables increased performance for AI and HPC
NVIDIA NGC	The NGC catalog provides a collection of GPU-optimized containers for AI and HPC
Slurm	A classic workload manager used to manage complex workloads in a multi-node, batch-style, compute environment
NVIDIA Run:AI	NVIDIA Run:ai accelerates AI operations with dynamic orchestration across the AI lifecycle. It maximizes GPU utilization, scales AI workloads efficiently, and integrates seamlessly into hybrid infrastructure with minimal overhead. Browse the documentation to install and monitor your environment, manage resources and organizations, run and scale AI workloads, and integrate NVIDIA Run:ai into your workflows using APIs.

Design Requirements

DGX SuperPOD is designed to minimize system bottlenecks throughout the tightly coupled configuration to provide the best performance and application scalability. Each subsystem has been thoughtfully designed to meet this goal. In addition, the overall design remains flexible so that SuperPOD can be tailored to better integrate into existing data centers.

System Design

DGX SuperPOD is optimized for a customers' particular workload of multi-node AI and HPC applications:

- A modular architecture based on Scalable Units (SUs) of 8 DGX Vera Rubin NVL72 rack-scale systems.
- > A fully tested system scales to two SUs, but larger deployments can be built based on customer requirements
 - A fully tested system scales for a single SU, seamlessly.
 - Rack-level integrated design which allows rapid installation and deployment of liquid cooled, high density compute racks.
- > Storage partner equipment that has been certified to work in a DGX SuperPOD environment.
 - Storage partner equipment that has been certified to work in DGX SuperPOD environments.
 - Full system support—including compute, storage, network, and software—is provided by NVIDIA Enterprise Support (NVEX).

Compute Fabric

- The compute fabric is rail-optimized to the top layer of the fabric.
- The compute fabric is a balanced, full-fat tree.
- The compute fabric is designed with current state-of-the-art, top performing, high-performance, low-latency network switches, and supports future generation of networking hardware.
- Managed XDR switches are used throughout the design to provide better management of the fabric.
- The fabric is designed to support the latest SHARP features.

Ethernet Fabric

- The compute fabric is connected to the BlueField DPUs on each compute tray, providing access to control plane, and high-performance storage.
- The compute fabric is a balanced design, with support for High Availability (HA) and provides sufficient bandwidth for the cluster for data ingestion and egress.
- It is independent of the compute fabric to maximize performance of both storage, control plane and application performance.
- The ethernet Fabric is shared among different network segments, namely:
 - Storage Network (High Speed Storage)
 - In-Band Management Network

The storage Network provides high bandwidth, low latency access to the shared high-performance storage. It also has the following characteristics:

- > Storage is provided on RDMA over Converged Ethernet to provide maximum performance and minimize CPU overhead.
- > It is flexible and can scale to meet specific capacity and bandwidth requirements.

In-Band Management Network has the following additional requirements:

- > The in-band management network fabric is Ethernet-based and is used for node provisioning, data movement, Internet access, and other services that must be accessible by the users.
- > The in-band management network connections for compute and management nodes operate at 400 Gbps and are running over a L3 wECMP routed connection together with the storage Network over the Ethernet Fabric for resiliency, which supports automatic failover between different ports.

Out-of-Band Management Network

The OOB management network connects all the base management controller (BMC) ports, the NVLink Switch trays, as well as other devices that should be physically isolated from users. The Switch Management Network is a subset of the Out-Of-Band Network that provides additional security and resiliency and provides direct ZTP capability from a single switch boot-strapping point.

Customers are encouraged to provide dedicated uplink handover point for the OOB network to allow dedicated access to the cluster independent of the main uplink.

Storage Requirements

The DGX SuperPOD compute architecture must be paired with a high-performance, balanced, storage system to maximize overall system performance. DGX SuperPOD is designed to use two separate storage systems, high-performance storage (HPS) and user storage, optimized for key operations of throughput, parallel I/O, as well as higher IOPS and metadata workloads.

High-Performance Storage

High-Performance Storage is provided via RDMA over Converged Ethernet v2 (RoCEv2) connected storage from a DGX SuperPOD certified storage partner, and is engineered and tested with the following attributes in mind:

- High-performance, resilient, POSIX-style file system optimized for multi-threaded read and write operations across multiple nodes.
- Certified for Grace-based systems.
- Native RoCE support
- Local system RAM for transparent caching of data.
- Leverage local flash device transparently for read and write caching.

The specific storage fabric topology, capacity, and components are determined by the DGX SuperPOD certified storage partner as part of the DGX SuperPOD design process.

User Storage

User Storage differs from High-Performance storage in that it exposes an NFS share on the in-band management fabric for multiple uses. It is typically used for “home directory” type usage (especially with clusters deployed with SLURM), administrative scratch space, and shared storage as needed by DGX SuperPOD components in a High Availability configuration (e.g., Mission Control), the requirement for log files collection, and system configuration files.

With that in mind, User Storage has the following recommended attributes:

- Designed for high metadata performance, IOPS, and key enterprise features such as log collection, data capacity. This is different than the HPS, which is optimized for parallel I/O and large capacity.
- Communicate over Ethernet, using NFS.
- 100GbE DR1 Connectivity at least.

User storage in a DGX SuperPOD is often satisfied with existing NFS servers already deployed, such that a new export is created and made accessible to the DGX SuperPOD’s in-band management network. User Storage is therefore not described in detail in this DGX SuperPOD reference architecture.

Network Fabrics

Building systems by scalable units (SUs) provides the most efficient designs. However, if a different node count is required due to budgetary constraints, data center constraints, or other needs, the fabric should be designed to support the full SU, including leaf switches and leaf-spine cables, and leave the portion of the fabric unused where these nodes would be located. This will ensure optimal traffic routing and ensure that performance is consistent across all portions of the fabric.

DGX SuperPOD configurations utilize five network segments, these segments are:

- Multi-node NVLink Fabric with NVLink 6
- Compute Network
- Storage Network
- In-Band Management Network
- Out-of-Band Management Network

These network segments are carried by four different physical fabrics:

- Multi-node NVLink Fabric (MN-NVL)
- Compute InfiniBand Fabric
- Ethernet Fabric
- Out-of-Band network

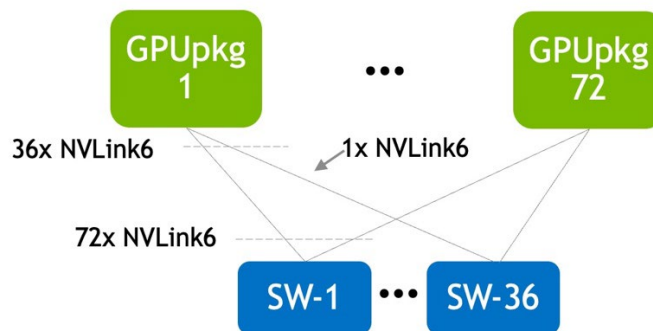
Each of these fabrics are discussed in this section.

Multi-node NVLink Fabric (NVL6)

Each DGX Vera Rubin NVL72 rack is built with 18 compute trays and 9 NVLink switch trays. Each NVLink switch tray is equipped with 2 NVLink switch chips and are responsible for the full-mesh connectivity between all 72 GPUs within the same DGX Vera Rubin NVL72 rack. Each Rubin GPU features 18 NVL6 links and has one dedicated NVL6 link connectivity to each one of the 18 switch chips, delivering a total bandwidth of 3.6 TB/s low latency bandwidth.

Figure 3 visualizes the topology of the Multi-node NVLink topology.

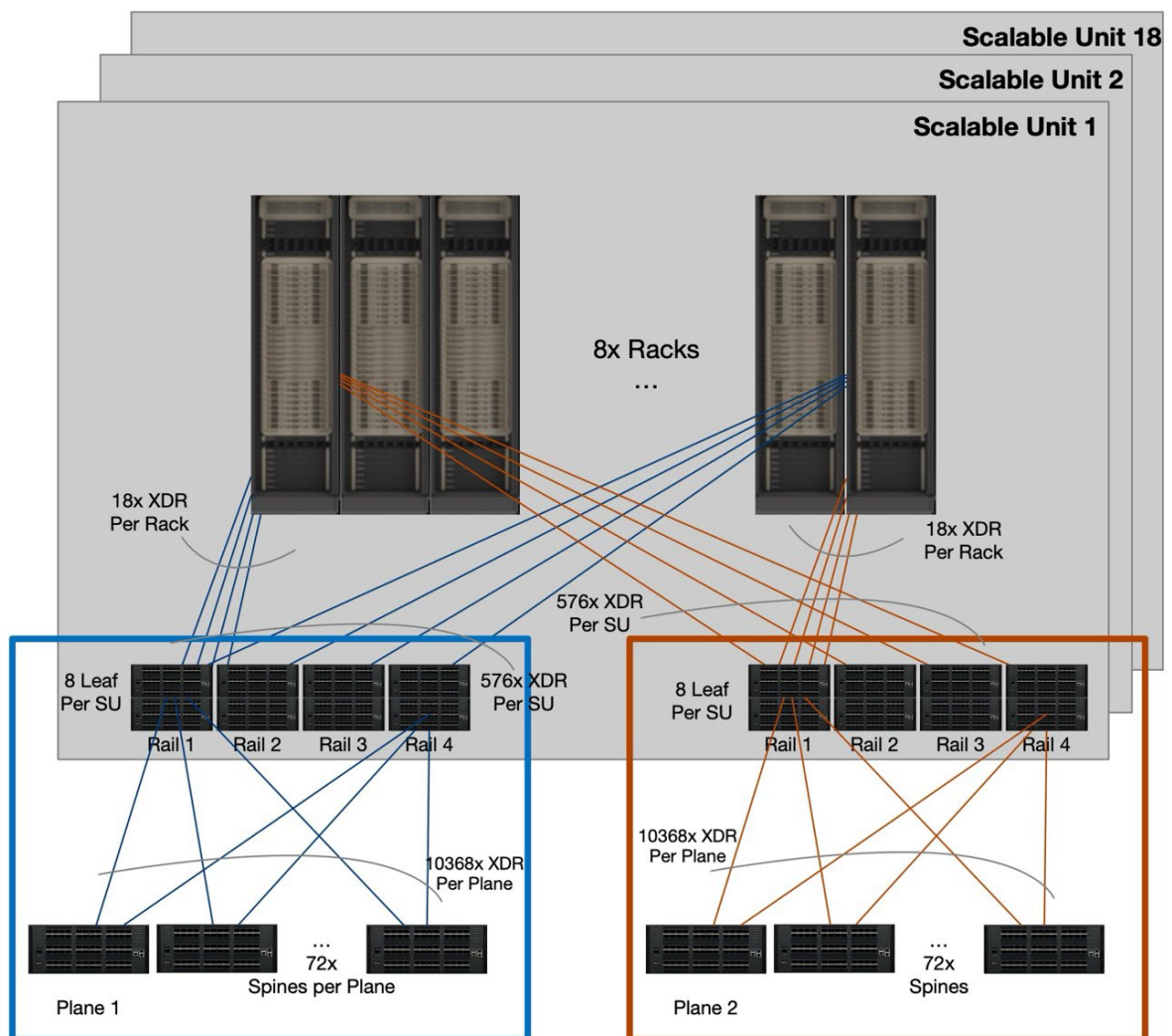
Figure 3: Multi-node NVLink6 Topology



Compute Fabric

Figure 4 shows the compute fabric layout for the full Vera Rubin NVL72 SuperPOD Scalable Unit (8 DGX Vera Rubin NVL72 systems). Each compute rack is rail-aligned. Traffic per rail of each compute tray is always one hop away from other compute trays in the same Scalable Unit. Traffic between different compute racks, or between different rails, traverses the spine layer. There are 4 rails per plane to connect to the compute.

Figure 4. Compute fabric for 18 SUs of 576 GPUs per SU Full DGX SuperPOD*



Thanks to the enhanced port radix of 144 port for Q3400-RD Quantum-X800 switches, we can scale the SuperPOD up to 18 SUs with a two-layer fabric. Each

SU contains a total of sixteen leaf switches distributed in 8 rails, with each rail having 2 leaf switches.

To leverage the 2 ConnectX-9 HCAs per GPU, two independent fabrics, named “planes” are used. The failure of a single plane can be tolerated as each GPU will still have a second path to another GPU in any location of the SuperPOD. Up to 16 SUs are aggregated in the 2-layer fabric – connecting to 10k GPUs in total.

Figure 4 shows the compute fabric design for DGX SuperPOD with Vera Rubin NVL72.

*Subject to Change

Figure 5. Q3400-RD InfiniBand Quantum-X800 switch

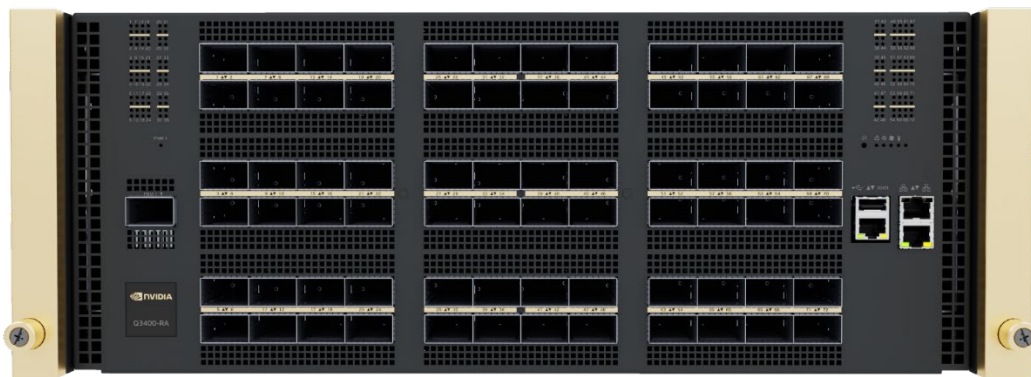


Table 3. Larger SuperPOD component counts

# GPUs	# SUs	# Planes	IB Leaf Switch		IB Spine Switch	
			Per Plane	Total	Per Plane	Total
1152	2	2	16	32	9	18
2304	4	2	32	64	18	36
4608	8	2	64	128	36	72
10368	18	2	144	288	72	144

Tenant Access Network (NS)

With DGX Vera Rubin NVL72, we continue to use the previously introduced Ethernet Fabric that was introduced with GB300 SuperPOD for storage and in-band network to enhance cost-efficiency while maintaining the high level of performance required by the storage network in a large-scale AI training cluster.

The storage and in-band fabric use SN5600 switches shown in Figure 6.

Figure 6. SN5600 switch



Figure 7 represents the scale-out design for Ethernet Fabric. The design follows a 2-layer design with spine-leaf switches. Each SU contains four leaf SN5600 switches. These switches are connected to the BlueField-4 B4240V dual-port card at 400Gbps for best performance and redundancy.

The 2-layer fabric with up to 24 spine switches supports DGX SuperPODs with up to 18 SUs (10k GPUs). Each of the ethernet features also a storage leaf side – with a configurable number of switches starting at 2 leaves. Finally, the Ethernet fabric is extended by a pair of management leaves to connect the persistent home storage, and the management nodes, as well as a pair of out-of-band-spines to connect the vast amount of SN2201 OOB switches.

For DGX Vera Rubin NVL72 SuperPOD scale-out, each compute leaf connects via 24 links at 400 Gbps, supporting a 1:3 blocking ratio. The storage leaves connect via 72 non-blocking 400 Gbps links. The management leaves are connected at 24x 400 Gbps as well to ensure best performance delivery for provisioning and home storage.

Figure 7. Storage and In-band Ethernet fabric logical design

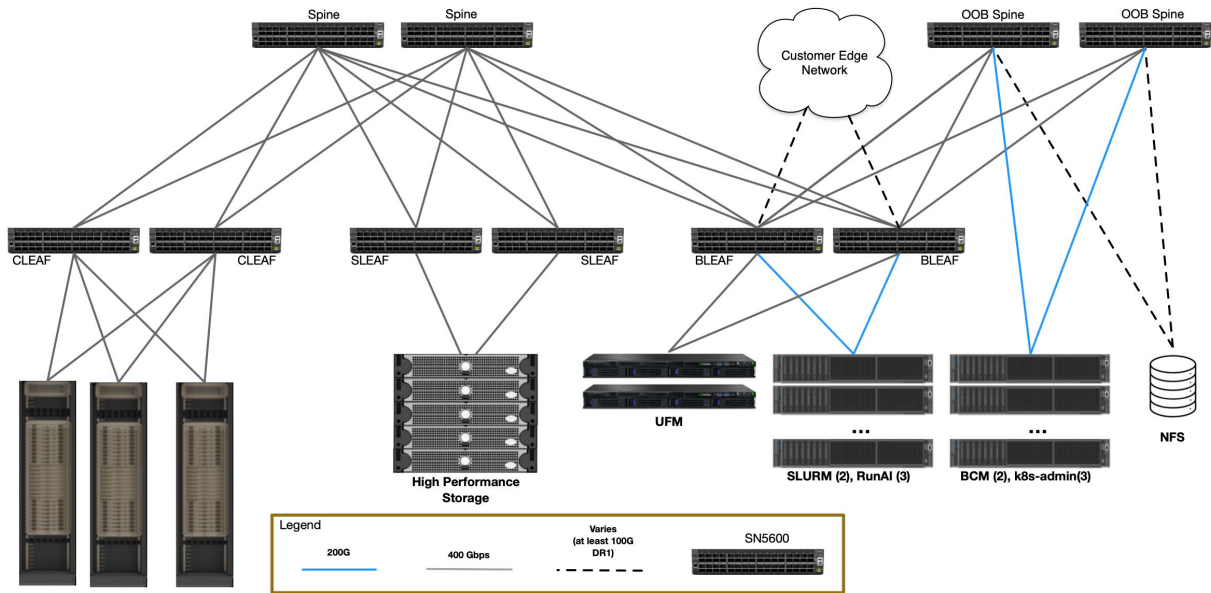


Table 4. Spine and Super Spine Switch Requirements for Scale Out

# GPUs	# SUs	Spines	Compute Leafs	Storage Leafs	Max. Supported Storage Bandwidth	Management Leafs and OOB Spines
1152	2	4	8	2	64x400GbE	2+2
2304	4	8	16	2	128x400GbE	2+2
4608	8	12	32	4	256x400GbE	2+2
10368	18	24	72	8	1024x400GbE	2+2

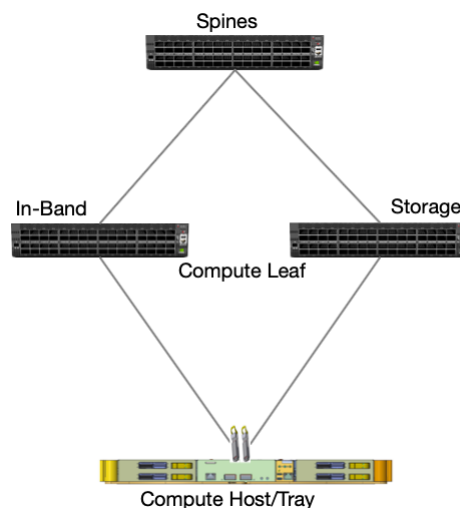
TAN - Network Segments

The Ethernet fabric on the SuperPOD is divided into the following segments:

- Storage Network
- In-band Network
- Out-of-Band Management Network

Figure 8 denotes a schematic design of ethernet fabric.

Figure 8. Network Segmentation diagram



The NVIDIA BlueField-4 DPU ports operate in NIC mode, with one interface connecting to a Compute Leaf (CLEAF) switch for in-band traffic and the second interface connecting to a separate CLEAF switch for storage traffic. These ports are configured as Layer 3 interfaces facing the leaf layer.

Within NVIDIA DGX SuperPOD architectures, the converged Ethernet fabric is deployed as a Layer 3 routed network. Standardizing on Layer 3 routing with EVPN delivers the high availability and efficiency required by modern AI data centers. This converged fabric employs an underlay network driven by Border Gateway Protocol (BGP) with Ethernet Virtual Private Network (EVPN) control-plane extensions and a VXLAN-based data plane. Both in-band and storage subnets terminate at the VXLAN Tunnel Endpoints (VTEPs) on the leaf switches.

NVIDIA Spectrum-4 (SN5600) switches provide full hardware EVPN support, ensuring optimal performance for both RoCEv2 and standard IP overlay traffic.

Storage and In-band Network

Storage and in-band traffic share a unified converged Ethernet fabric, achieving logical isolation via separate VXLAN Network Identifiers (VNIs) mapped to VXLAN Tunnel Endpoints (VTEPs) terminated at the Compute Leaf switches.

Storage Optimization

High-performance storage access relies on RDMA over Converged Ethernet version 2 (RoCEv2) to deliver line-rate throughput and ultra-low latency. To meet the stringent lossless delivery requirements of RoCEv2, the fabric implements end-to-end flow and congestion control mechanisms.

- Priority-Based Flow Control (PFC): Eliminates packet drops caused by buffer exhaustion on designated storage traffic classes.
- Explicit Congestion Notification (ECN): Proactively signals congestion to endpoints, triggering dynamic rate-throttling before buffers saturate.

In-Band Capabilities

The in-band network functions as the cluster's operational backbone, providing:

- Cluster Service: Facilitates node-to-node communication among all core infrastructure and control-plane services.
- Storage Access: Serves as the transport path for lower-speed, capacity-oriented Network File System (NFS) storage tiers.
- Border Connectivity: Routes traffic between internal management stacks (NVIDIA Mission Control, Base Command Manager, Slurm, and Kubernetes) and other services outside of the cluster such as the code repositories, and data sources.
- Client Connectivity: Enables secure administrative and end-user access to Slurm login/head nodes and Kubernetes ingress endpoints.

Out-of-Band Management Network

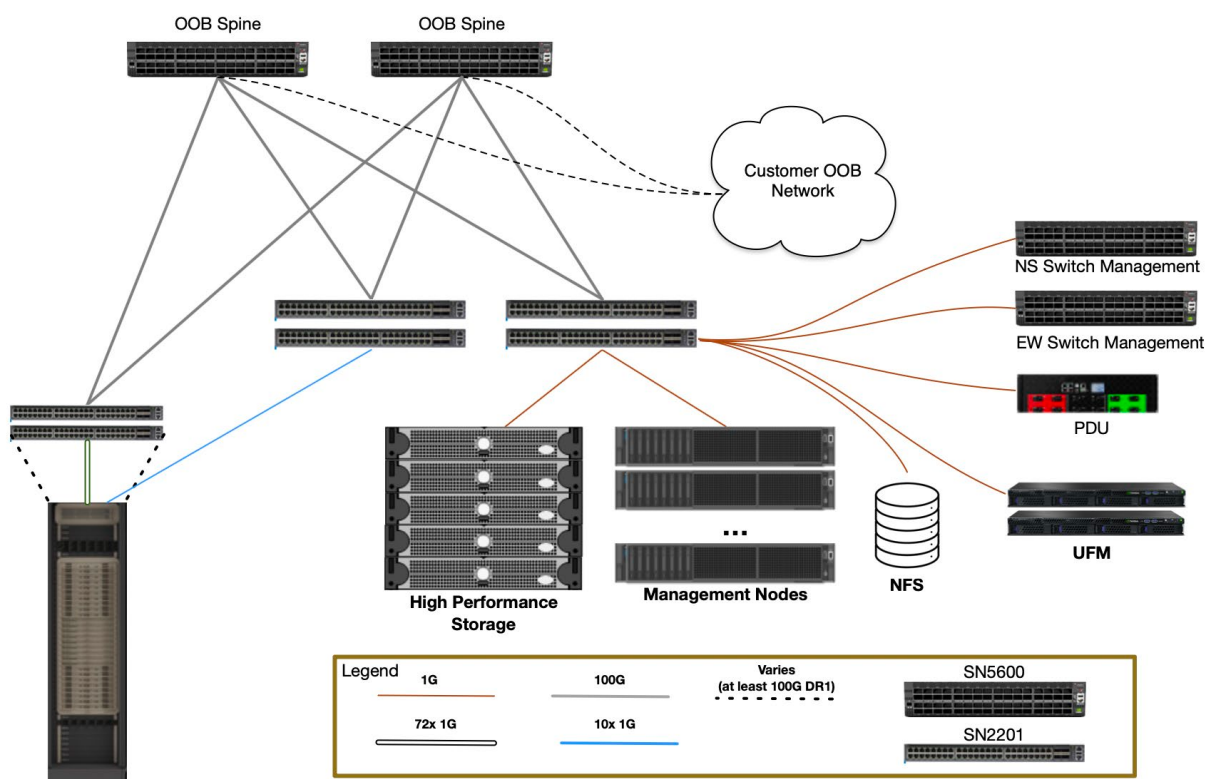
The out-of-band management network (OOB) carries all IPMI related control traffic. It connects the management ports of all devices including DGX Vera Rubin NVL72 Racks (both BF4 BMC and BMC), Infiniband and Ethernet switches, management servers, storage appliances, rack PDUs and power shelves, and all other devices that are part of the SuperPOD and supports out-of-band management feature. These are separated onto their own network, with dedicated, separated uplink for a parallel way of getting into the SuperPOD in case of disaster recovery or for management purposes.

The OOB network is physically rolled up into the management leaf (MLEAF). The OOB management network use SN2201 switches as top-of-the-rack (TOR) switches (shown in Figure 9), and SN5600D switches as spine switches.

Figure 9. SN2201 switch



Figure 10: Schematics for the Out-of-band Network



High Performance Storage Architecture

Data, lots of data, is the key to the development of accurate deep learning (DL) models. Data volume continues to grow exponentially, and data used to train individual models continues to grow as well. Data format, not just volume can play a key factor in the rate at which data is accessed so storage system performance must scale commensurately.

The key I/O operation in DL training is re-read. It is not just that data is read, but it must be reused repeatedly due to the iterative nature of DL training. Pure read performance still is important as some model types can train in a fraction of an epoch (ex: some recommender models) and inference of existing can be highly I/O intensive, much more so than training. Write performance can also be important. As DL models grow and time-to-train, writing checkpoints is necessary for fault tolerance. The size of checkpoint files can be terabytes in size and while not written frequently are typically written synchronously than blocks and block forward progress of DL models.

Ideally, data is cached during the first read of the dataset, so data does not have to be retrieved across the network. Shared filesystems typically use RAM as the first layer of cache. Reading files from cache can be an order of magnitude faster than from remote storage. In addition, the DGX Vera Rubin NVL72 system provides local NVMe storage that can also be used for caching or staging data.

DGX SuperPOD is designed to support all workloads, but the storage performance required to maximize training performance can vary depending on the type of model and dataset. The guidelines in Table 5 and 6 are provided to help determine the I/O levels required for different types of models.

Table 5. Storage performance requirements

Level	Work Description	Dataset Size
Standard	Multiple concurrent LLM or fine-tuning training jobs and periodic checkpoints, where the compute requirements dominate the data I/O requirements significantly.	Most datasets can fit within the local compute systems' memory cache during training. The datasets are single modality, and models have millions of parameters.
Enhanced	Multiple concurrent multimodal training jobs and periodic checkpoints, where the data I/O performance is an important factor for end-to-end training time.	Datasets are too large to fit into local compute systems' memory cache requiring more I/O during training, not enough to obviate the need for frequent I/O. The datasets have multiple modalities and models have billions (or higher) of parameters.

Table 6. Guidelines for storage performance

Performance Characteristic	Standard (GBps)	Enhanced (GBps)
Single SU aggregate system read	90	280
Single SU aggregate system write	45	140
4 SU aggregate system read	360	1,120
4 SU aggregate system write	180	560

High-speed storage provides a shared view of an organization's data to all nodes. It must be optimized for small, random I/O patterns, and provide high peak node performance and high aggregate filesystem performance to meet the variety of workloads an organization may encounter. High-speed storage should support both efficient multi-threaded reads and writes from a single system, but most DL workloads will be read-dominant.

Use cases in automotive and other computer vision-related tasks, where high-resolution images are used for training (and in some cases are uncompressed) involve datasets that easily exceed 30 TB in size. In these cases, 4 GBps per GPU of read performance is required.

While NLP and LLM cases often do not require as much read performance for training, peak performance for reads and writes are needed for creating and reading checkpoint files. This is a synchronous operation, and training stops during this phase. If you are looking for best end-to-end training performance, do not ignore I/O operations for checkpoints. Consider at least ½ of the read performance as recommended write performance for LLM and large model use cases.

The preceding metrics assume a variety of workloads, datasets, and need for training locally and directly from the high-speed storage system. It is best to characterize workloads and organizational needs before finalizing performance and capacity requirements.

Management Servers

For DGX SuperPOD with DGX Rubin NVL72 systems, we provide two means of cluster access for end users (AI practitioners):

- > With SLURM Workload Manager
- > With Kubernetes and NVIDIA Run:ai.

For details on control plane architecture refer [NVIDIA Mission Control Architecture](#)

The control plane nodes shall have the following minimum configuration requirement:

- > All ARM for DGX SuperPOD; 2x x86 (amd64) based CPU, each 32-core minimum
- > 1TB of RAM
- > 2x NVMe 7.68TB for local storage
- > 2x 960GB SSD (RAID 1) for OS
- > 1x **dual port** BlueField 3 or BlueField 4 DPU
- > 2x 1GbE RJ45 Ports for PTP or additional low-speed networking connectivity
- > 1GbE RJ45 BMC / IPMI / Redfish, supporting remote media feature.

DGX SuperPOD Software

DGX SuperPOD is an integrated hardware and software solution for building large-scale AI infrastructure. NVIDIA Mission Control is an integrated AI factory management platform designed to simplify operations, reduce downtime, and accelerate model development for enterprise AI infrastructure. It combines NVIDIA's best operational practices and AI cluster automation into a single control plane.

For further details, please refer to the NVIDIA Mission Control docs at [NVIDIA Mission Control](#)

Summary

NVIDIA DGX SuperPOD with NVIDIA DGX Vera Rubin NVL72 system is the next generation of data center scale architecture to meet the demanding and growing needs of AI training and inferencing. This reference architecture document for DGX SuperPOD represents the architecture used by NVIDIA for our own AI model and HPC research and development. DGX SuperPOD continues to build upon its high-performance roots to enable training of the largest NLP models, support the expansive needs of training models for automotive applications, and scaling-up recommender models for greater accuracy and faster turn-around-time. DGX SuperPOD represents a complete system of not just hardware but all the necessary software to accelerate time-to-deployment, streamline system management, proactively identify system issues. The combination of all these components keeps systems running reliably, with maximum performance, and enables users to push the bounds of state-of-the-art. The platform is designed to both support the workloads of today and grow to support tomorrow's applications. These are representatives of the configuration and must be finalized based on actual design.

Your NVIDIA representative can work with you to finalize the exact components list.

Further information on the detailed requirement for DGX SuperPOD, such as the datacenter design guide, detailed software architecture for NVIDIA mission control, and relevant security information, are also available to customers through NVIDIA representative.

NVIDIA Reference Architecture Notice and Disclaimers

NVIDIA Reference Architecture documentation is NVIDIA confidential information and embodies substantial intellectual property and engineering know-how. NVIDIA grants you limited permissions to use NVIDIA Reference Architecture to create NVIDIA platform-enabled data center clusters with NVIDIA technologies. You may not use the NVIDIA Reference Architecture to develop competing products or technologies or assisting a third party in such activities. Except as expressly stated in this Notice & Disclaimer statement, no other license or right is granted to you by implication, estoppel or otherwise. NVIDIA objects to and rejects any additional or different terms you may propose.

NVIDIA Reference Architecture is subject to change without notice. NVIDIA technologies may require enabled hardware, software, or service activation. Your costs and results may vary. This document is not a commitment to develop, release, or deliver any technology, code, or functionality. NVIDIA reserves the right to make corrections, modifications, enhancements, improvements, and any other changes to this document, at any time without notice. It is your sole responsibility to evaluate and determine the applicability of any information contained in this document and ensure that the product is suitable and fit for the application planned by you. Information published by NVIDIA regarding third-party products or services does not constitute a license from NVIDIA to use such products or services or a warranty or endorsement thereof. Use of such information may require a license from a third party or a license from NVIDIA.

NVIDIA products or services are subject to the applicable NVIDIA product or sales terms and conditions that are provided in conjunction with the relevant NVIDIA product or service. NVIDIA's provision of these resources does not expand or otherwise alter NVIDIA's applicable warranties or warranty disclaimers for the NVIDIA products or services.

You will not use the NVIDIA Reference Architecture or NVIDIA's confidential information to: (i) identify or support an assertion or potential assertion of any intellectual property rights (including patent, copyright, or trade secret); or (ii) prepare, file, amend, or prosecute any patent applications. You agree to grant NVIDIA a limited, worldwide, nonexclusive, perpetual, irrevocable, non-sublicensable, nontransferable, royalty-free, fully-paid license to any patent claim that you may file after accessing NVIDIA Reference Architecture that covers the subject matter disclosed herein.

Warranty Disclaimer. NVIDIA Reference Architecture materials disclosed are provided "AS IS". To the maximum extent permitted by applicable law, NVIDIA disclaims all warranties and representations of any kind, whether express, implied, or statutory, relating to or arising under this agreement, including, without limitation, the warranties of title, noninfringement, merchantability, fitness for a particular purpose, usage of trade, and course of dealing.

Limitation of Liability. To the maximum extent permitted by applicable law, in no event will NVIDIA be liable for any (i) indirect, punitive, special, incidental, or consequential damages, or (ii) damages for the (a) cost of procuring substitute goods or (b) loss of profits, revenues, use, data, or goodwill arising out of or related to this document, whether based on breach of contract, tort (including negligence), strict liability, or otherwise, and even if NVIDIA has been advised of the possibility of such damages and even if a your remedies fail their essential purpose. Additionally, to the maximum extent permitted by applicable law, NVIDIA's total cumulative aggregate liability for any and all liabilities, obligations, or claims arising out of or related to this document will not exceed one-hundred U.S. dollars (US\$100).

NVIDIA, NVIDIA HGX, NVIDIA DGX SuperPOD, NVIDIA DGX Cloud, NVIDIA ConnectX, NVIDIA BlueField, NVIDIA Spectrum, NVIDIA NVLink, NVIDIA Base Command, NVIDIA CUDA, NVIDIA Magnum IO, NVIDIA NetQ, and the NVIDIA logo are trademarks and/or registered trademarks of NVIDIA Corporation in the U.S. and other countries. Other company and product names may be trademarks of the respective companies with which they are associated.

© 2024 NVIDIA Corporation & Affiliates. All rights reserved.