

Lenovo Hybrid AI 285 with Cisco Networking

Product Guide

This document is meant to be used in tandem with the [Hybrid AI 285 platform guide](#). It describes the hardware architecture changes required to leverage Cisco networking hardware and the Cisco Nexus Dashboard within the Hybrid AI 285 Platform.

Lenovo Hybrid AI 285 is a platform that enables enterprises of all sizes to quickly deploy hybrid AI factory infrastructure, supporting Enterprise AI use cases as either a new, greenfield environment or an extension of their existing IT infrastructure. This document will often send the user to the base [Hybrid AI 285 platform guide](#) as its main purpose is to show the main differences required to implement Cisco networking and Nexus dashboard.

The offering is based on the NVIDIA 2-8-5 PCIe-optimized configuration — 2x CPUs, 8x GPUs, and 5x network adapters — and is ideally suited for medium (per GPU) to large (per node) Inference use cases, and small-to-large model training or fine-tuning, depending on chosen scale. It combines market leading Lenovo ThinkSystem GPU-rich servers with NVIDIA Hopper or Blackwell GPUs, Cisco networking and enables the use of the NVIDIA AI Enterprise software stack with NVIDIA Blueprints.

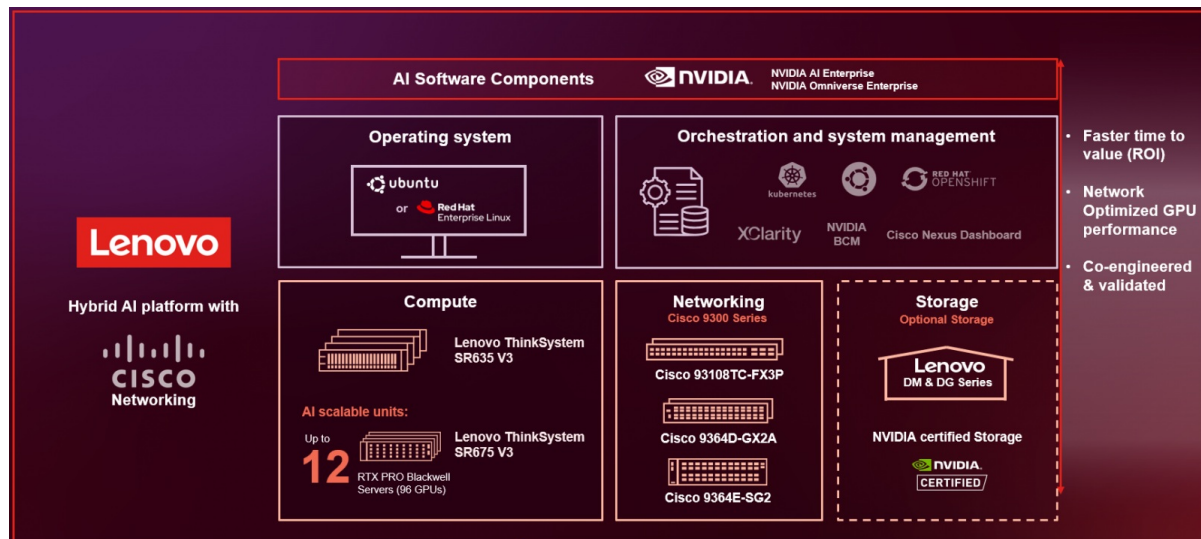


Figure 1. Lenovo Hybrid AI 285 platform overview with Cisco

Did you know?

The same team of HPC and AI experts that created the Lenovo EveryScale OVX solution, as deployed for NVIDIA Omniverse Cloud, brings the Lenovo Hybrid AI 285 with Cisco networking to market.

Following their excellent experience with Lenovo on Omniverse, NVIDIA has once again chosen Lenovo technology as the foundation for the development and test of their [NVIDIA AI Enterprise Reference Architecture](#) (ERA).

Overview

The Lenovo Hybrid AI 285 Platform with Cisco Networking scales from a **Starter Kit** environment with between 4-32 PCIe GPUs to a **Scalable Unit Deployment (SU)** with four servers and 32 GPUs in each SU and up to 3 Scalable Units with 12 Servers and 96 GPUs. See figure below for a sizing overview.

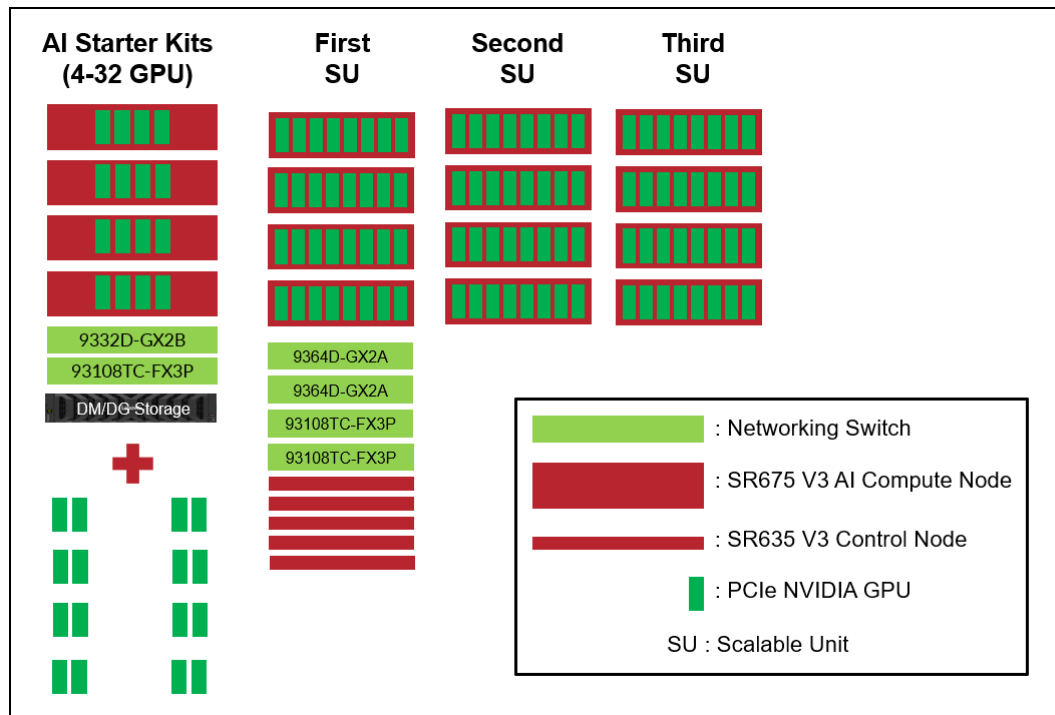


Figure 2. Lenovo Hybrid AI 285 with Cisco Networking scaling from Starter Kit to 96 GPUs

The figure below shows the networking architecture of the platform deployed with 96 GPUs.

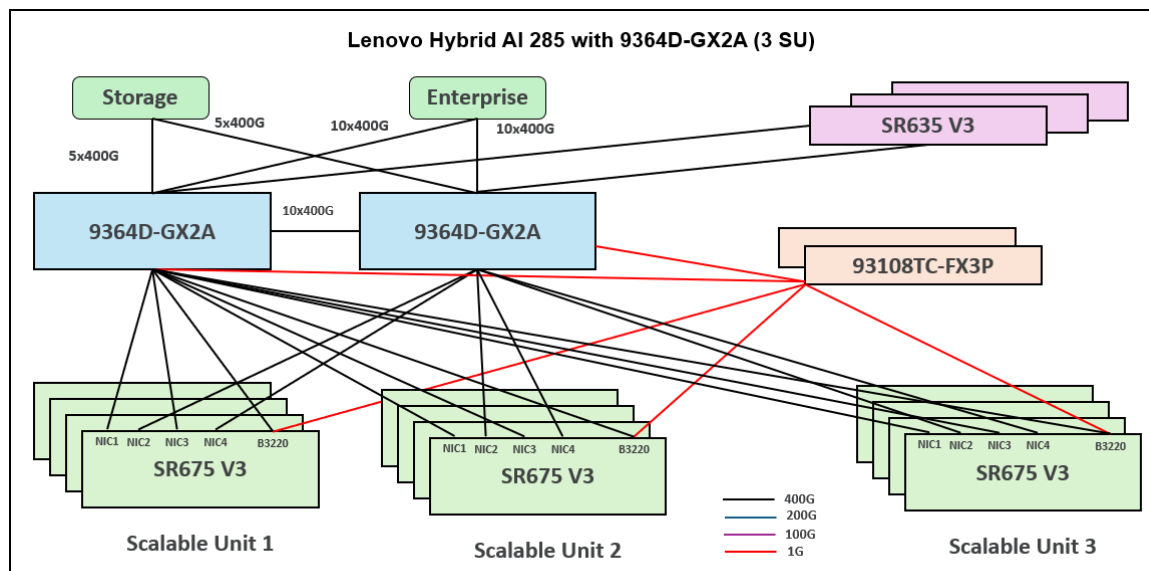


Figure 3. Lenovo Hybrid AI 285 with Cisco Networking platform with 3 Scalable Units

Components

The main hardware components of Lenovo Hybrid AI platforms are Compute nodes and the Networking infrastructure. As an integrated solution they can come together in either a Lenovo EveryScale Rack (Machine Type 1410) or Lenovo EveryScale Client Site Integration Kit (Machine Type 7X74).

Topics in this section:

- [AI Compute Node – SR675 V3](#)
- [Service Nodes – SR635 V3](#)
- [Cisco Networking](#)
- [Lenovo EveryScale Solution](#)

AI Compute Node – SR675 V3

The AI Compute Node leverages the [Lenovo ThinkSystem SR675 V3](#) GPU-rich server.

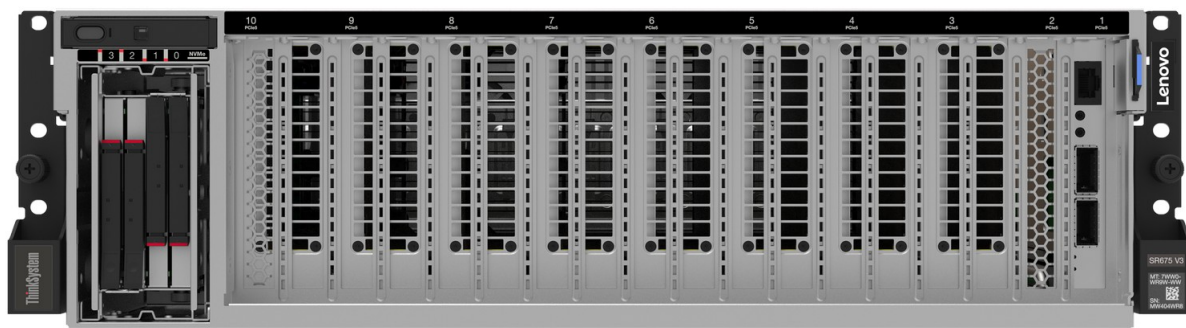


Figure 4. Lenovo ThinkSystem SR675 V3 in 8DW PCIe Setup

The SR675 V3 is a 2-socket 5th Gen AMD EPYC 9005 server supporting up to 8 PCIe DW GPUs with up to 5 network adapters in a 3U rack server chassis. This makes it the ideal choice for NVIDIA's 2-8-5 configuration requirement.

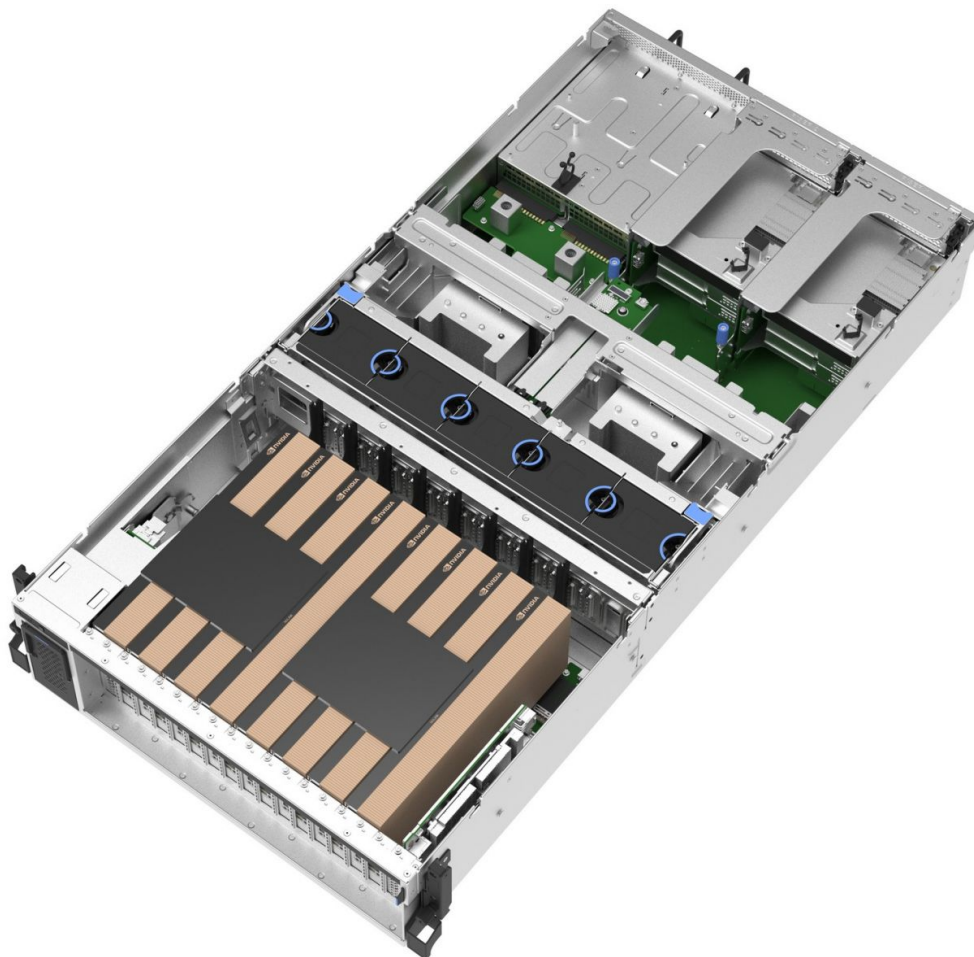


Figure 5. Lenovo ThinkSystem SR675 V3 in 8DW PCIe Setup

The SR675 V3 is configured as the 285 AI compute node as shown in the figure below.

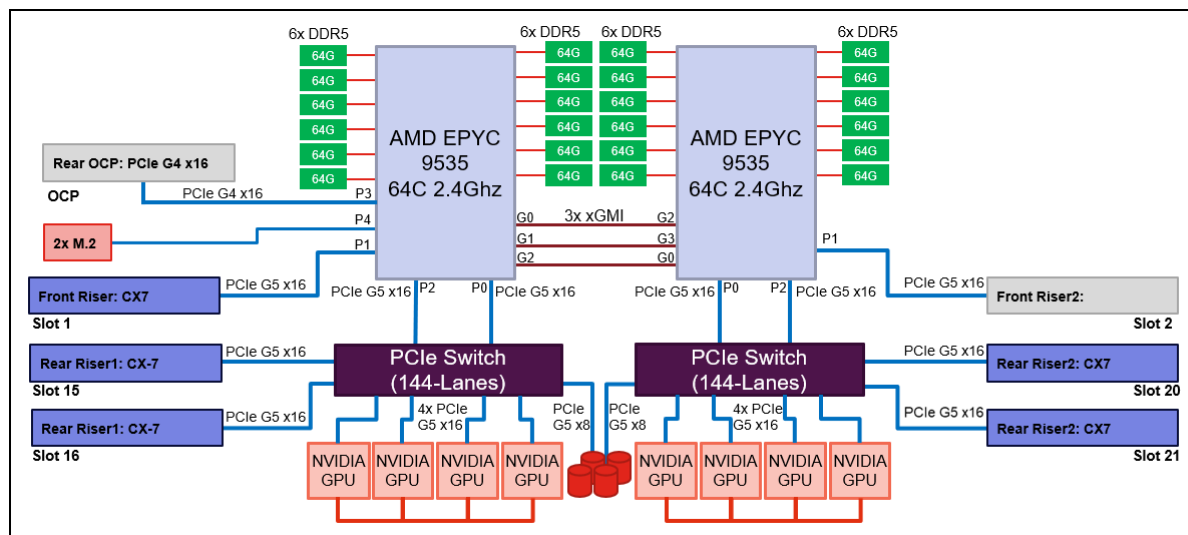


Figure 6. AI Compute Node Block Diagram

The AI Compute node is configured with two **AMD EPYC 9535 64 Core 2.4 GHz** processors with an all-core boost frequency of 3.5GHz. Besides providing consistently more than 2GHz frequency this ensures that with 7 Multi Instance GPUs (MIG) on 8 physical GPUs there are 2 Cores available per MIG plus a few additional Cores for Operating System and other operations.

With 12 memory channels per processor socket the AMD based server provides superior memory bandwidth versus competing Intel-based platforms ensuring highest performance. Leveraging 64GB 6400MHz Memory DIMMs for a total of 1.5TB of main memory providing 192GB memory per GPU or a minimum of 1.5X the H200 NVL GPU memory.

The GPUs are connected to the CPUs via two PCIe Gen5 switches, each supporting up to four GPUs. With the **NVIDIA H200 NVL PCIe GPU**, the four GPUs are additionally interconnected through an NVLink bridge, creating a unified memory space. In an entry configuration with two GPUs per PCIe switch, the ThinkSystem SR675 V3 uniquely supports connecting all four GPUs with an NVLink bridge for maximized shared memory, thereby accommodating larger inference models, rather than limiting the bridge to two GPUs. With the RTX PRO 6000 Blackwell Server Edition, no NVLink bridge is applicable, same applies to configurations with the L40S.

Key difference from the base platform : For the 2-8-5 architecture with Cisco networking the AI Compute node leverages NVIDIA CX-7s for both the East-West and the North South communication. This is a key difference in the AI compute node configuration compared to the base 285 platform with NVIDIA networking. This is done primarily because, as of now, Cisco switching does not work with NVIDIA's dynamic load balancing technology within Spectrum-X, leveraged by their Bluefield-3 cards. This will change in future updates of this document as that technology is onboarded by Cisco and Lenovo brings in the new Silicon One 800GbE switch.

The Ethernet adapters for the Compute (East-West) Network are directly connected to the GPUs via PCIe switches minimizing latency and enabling NVIDIA GPUDirect and GPUDirect Storage operation. For pure Inference workload they are optional, but for training and fine-tuning operation they should provide at least 200Gb/s per GPU.

Finally, the system is completed by local storage with two 960GB Read Intensive M.2 in RAID1 configuration for the operating system and four 3.84TB Read Intensive E3.S drives for local application data.

GPU selection

The Hybrid AI 285 platform is designed to handle any of NVIDIA's DW PCIe form factor GPUs including the new RTX PRO 6000 Blackwell Server Edition, the H200 NVL, L40S and the H100 NVL.

- **NVIDIA H200 NVL**

The NVIDIA H200 NVL is a powerful GPU designed to accelerate both generative AI and high-performance computing (HPC) workloads. It boasts a massive 141GB of HBM3e memory, which is nearly double the capacity of its predecessor, the H100. This increased memory, coupled with a 4.8 terabytes per second (TB/s) memory bandwidth, enables the H200 NVL to handle larger and more complex AI models, like large language models (LLMs), with significantly improved performance. In addition, the H200 NVL is built with energy efficiency in mind, offering increased performance within a similar power profile as the H100, making it a cost-effective and environmentally conscious choice for businesses and researchers.

NVIDIA provides a 5-year license to NVIDIA AI Enterprise free-of-charge bundled with NVIDIA H200 NVL GPUs.

- **NVIDIA RTX PRO 6000 Blackwell Server Edition**

Built on the groundbreaking NVIDIA Blackwell architecture, the NVIDIA RTX PRO™ 6000 Blackwell Server Edition delivers a powerful combination of AI and visual computing capabilities to accelerate enterprise data center workloads. Equipped with 96GB of ultra- fast GDDR7 memory, the NVIDIA RTX PRO 6000 Blackwell provides unparalleled performance and flexibility to accelerate a broad range of use cases- from agentic AI, physical AI, and scientific computing to rendering, 3D graphics, and video.

Configuration

The following table lists the configuration of the AI Compute Node with H200 NVL GPUs.

Table 1. AI Compute Node

Part Number	Product description	Quantity per system
7D9RCTOLWW	ThinkSystem SR675 V3	
BR7F	ThinkSystem SR675 V3 8DW PCIe GPU Base	1
C3EF	ThinkSystem SR675 V3 System Board v2	1
C2AL	ThinkSystem AMD EPYC 9535 64C 300W 2.4GHz Processor	2
C0CK	ThinkSystem 64GB TruDDR5 6400MHz (2Rx4) RDIMM-A	24
BR7S	ThinkSystem SR675 V3 Switched 4x16 PCIe DW GPU Direct RDMA Riser	2
C3V3	ThinkSystem NVIDIA H200 NVL 141GB PCIe GPU Gen5 Passive GPU	8
C3V0	ThinkSystem NVIDIA 4-way bridge for H200 NVL	2
BR7H	ThinkSystem SR675 V3 2x16 PCIe Front IO Riser	1
C2RK	ThinkSystem SR675 V3 2 x16 Switch Cabled PCIe Rear IO Riser	2
BQBN	ThinkSystem NVIDIA ConnectX-7 NDR200/200GbE QSFP112 2-port PCIe Gen5 x16 Adapter	
BM8X	ThinkSystem M.2 SATA/x4 NVMe 2-Bay Adapter	1
BT7P	ThinkSystem Raid 540-8i for M.2/7MM NVMe boot Enablement	1
BXMH	ThinkSystem M.2 PM9A3 960GB Read Intensive NVMe PCIe 4.0 x4 NHS SSD	2
BTMB	ThinkSystem 1x4 E3.S Backplane	1
C1AB	ThinkSystem E3.S PM9D3a 3.84TB Read Intensive NVMe PCIe 5.0 x4 HS SSD	2
BK1E	ThinkSystem SR670 V2/ SR675 V3 OCP Enablement Kit	1
C5WW	ThinkSystem SR675 V3 Dual Rotor System High Performance Fan	5
BFD6	ThinkSystem SR670 V2/ SR675 V3 Power Mezzanine Board	1
BE0D	N+1 Redundancy With Over-Subscription	1
BKTJ	ThinkSystem 2600W 230V Titanium Hot-Swap Gen2 Power Supply	4
6252	2.5m, 16A/100-250V, C19 to C20 Jumper Cord	4
C3KA	ThinkSystem SR670 V2/SR675 V3 Heavy Systems Toolless Slide Rail Kit	1
BFNU	ThinkSystem SR670 V2/ SR675 V3 Intrusion Cable	1
BR7U	ThinkSystem SR675 V3 Root of Trust Module	1
BFTH	ThinkSystem SR670 V2/ SR675 V3 Front Operator Panel ASM	1
5PS7B09631	5Yr Premier NBD Resp + KYD SR675 V3	1

Service Nodes – SR635 V3

When deploying the Hybrid AI 285 platform in a sizing beyond beyond 2 AI compute nodes additional service nodes are recommended to manage the overall AI cluster environment.

Key difference from the base platform: Because the architecture for this platform does not yet leverage Spectrum-X, there is no need for Bluefields within the service nodes. For this reason the customer can leverage the lower cost SR635 V3 instead of the SR655 V3.

Two **Management Nodes** provide a high-availability for the System Management and Monitoring provided through NVIDIA Base Command Manager (BMC) as described further in the AI Software Stack chapter.

For the Container operations three **Scheduling Nodes** build the Kubernetes control plane providing redundant operations and quorum capability.



Figure 7. Lenovo ThinkSystem SR635 V3

The [Lenovo ThinkSystem SR635 V3](#) is an optimal choice for a homogeneous host environment, featuring a single socket AMD EPYC 9335 with 32 cores operating at 3.0 GHz base with an all-core boost frequency of 4.0GHz. The system is fully equipped with twelve 32GB 6400MHz Memory DIMMs, two 960GB Read Intensive M.2 drives in RAID1 configuration for the operating system, and two 3.84TB Read Intensive U.2 drives for local data storage. Additionally, it includes a **NVIDIA dual port CX7** adapter to connect the Service Nodes to the Converged Network.

Configuration

The following table lists the configuration of the Service Nodes.

Table 2. Service Nodes

Part Number	Product description	Quantity per system
7D9GCTO1WW	Server : ThinkSystem SR635 V3 - 3yr Warranty	
BLK4	ThinkSystem V3 1U 10x2.5" Chassis	1
BVGL	Data Center Environment 30 Degree Celsius / 86 Degree Fahrenheit	1
C2AQ	ThinkSystem AMD EPYC 9335 32C 210W 3.0GHz Processor	1
BQ26	ThinkSystem SR645 V3/SR635 V3 1U High Performance Heatsink	1
C1PL	ThinkSystem 32GB TruDDR5 6400MHz (1Rx4) RDIMM-A	12
BC4V	Non RAID NVMe	1
C0ZU	ThinkSystem 2.5" U.2 VA 3.84TB Read Intensive NVMe PCIe 5.0 x4 HS SSD	2
BPC9	ThinkSystem 1U 4x 2.5" NVMe Gen 4 Backplane	1
B5XJ	ThinkSystem M.2 SATA/NVMe 2-Bay Adapter	1
BTTY	M.2 NVMe	1
BKSR	ThinkSystem M.2 7450 PRO 960GB Read Intensive NVMe PCIe 4.0 x4 NHS SSD	2
BQBN	ThinkSystem NVIDIA ConnectX-7 NDR200/200GbE QSFP112 2-port PCIe Gen5 x16 Adapter	1
BLK7	ThinkSystem SR635 V3/SR645 V3 x16 PCIe Gen5 Riser 1	1
BLK9	ThinkSystem V3 1U MS LP+LP BF Riser Cage	1
BNFG	ThinkSystem 750W 230V/115V Platinum Hot-Swap Gen2 Power Supply v3	2
BH9M	ThinkSystem V3 1U Performance Fan Option Kit v2	7
BLKD	ThinkSystem 1U V3 10x2.5" Media Bay w/ Ext. Diagnostics Port	1
7Q01CTS2WW	5Yr Premier NBD Resp + KYD SR635 V3	1

Cisco Networking

The default setup of the Lenovo Hybrid AI 285 platform leverages Cisco Networking with the Nexus 9364D-GX2A for the Converged and Compute Network and the Nexus 9300-FX3 for the Management Network.

Cisco Nexus 9300-GX2A Series Switches

The Cisco Nexus 9364D-GX2A is a 2-rack-unit (2RU) switch that supports 25.6 Tbps of bandwidth and 8.35 bpps across 64 fixed 400G QSFP-DD ports and 2 fixed 1/10G SFP+ ports. QSFP-DD ports also support native 200G (QSFP56), 100G (QSFP28), and 40G (QSFP+). Each port can also support 4 x 10G, 4 x 25G, 4 x 50G, 4 x 100G, and 2 x 200G breakouts.

It supports flexible configurations, including 128 ports of 200GbE or 256 ports of 100/50/25/10GE ports accommodating diverse AI/ML cluster requirements.



Figure 8. Nexus 9364D-GX2A

The **Converged (North-South) Network** handles storage and in-band management, linking the Enterprise IT environment to the Agentic AI platform. Built on Ethernet with RDMA over Converged Ethernet (RoCE), it supports current and new cloud and storage services as outlined in the AI Compute node configuration.

In addition to providing access to the AI agents and functions of the AI platform, this connection is utilized for all data ingestion from the Enterprise IT data during indexing and embedding into the Retrieval-Augmented Generation (RAG) process. It is also used for data retrieval during AI operations.

The Storage connectivity is exactly half that and described in the Storage Connectivity chapter.

The **Compute (East-West) Network** facilitates application communication between the GPUs across the Compute nodes of the AI platform. It is designed to achieve minimal latency and maximal performance using a rail-optimized, fully non-blocking fat tree topology with Cisco Nexus 9300 series switches.

Cisco Nexus 9000 series data Center switches deliver purpose-built networking solutions designed specifically to address these challenges, providing the foundation for scalable, high-performance AI infrastructures that accelerate time-to-value while maintaining operational efficiency and security. Built on Cisco's custom Cloud Scale and Silicon One ASICs, these switches provide a comprehensive solution for AI-ready data centers.

Tip: In a pure Inference use case, the Compute Network is typically not necessary, but for training and fine-tuning operations it is a crucial component of the solution.

For configurations of up to five Scalable Units, the Compute and Converged Network are integrated utilizing the same switches. When deploying more than five units, it is necessary to separate the fabric.

The following table lists the configuration of the Cisco Nexus 9364D-GX2A.

Table 3. Cisco Nexus 9364D-GX2A configuration

Part Number	Description	Quantity
7DLKCTO1WW	Cisco Nexus 9300-GX2 Series Switch (N9K-C9364D-GX2A)	
C5P0	Cisco Nexus 9300-GX2 Series Switch (N9K-C9364D-GX2A)	2
C6FK	Mode selection between ACI and NXOS (MODE-NXOS)	2
6252	2.5m, 16A/100-250V, C19 to C20 Jumper Cord	4
C1P1TN9300XF2-5Y	5 Years (60 months) Cisco software Premier license	2

Cisco Nexus 9300-FX3 Series Switch

The Cisco Nexus 93108TC-FX3P is a high-performance, fixed-port switch designed for modern data centers. It features 48 ports of 100M/1/2.5/5/10GBASE-T, providing flexible connectivity options for various network configurations. Additionally, it includes 6 uplink ports that support 40/100 Gigabit Ethernet QSFP28, ensuring high-speed data transfer and scalability.

Built on Cisco's CloudScale technology, the 93108TC-FX3P delivers exceptional performance with a bandwidth capacity of 2.16 Tbps and the ability to handle up to 1.2 billion packets per second (Bpps).

This switch also supports advanced features such as comprehensive security, telemetry, and automation capabilities, which are essential for efficient network management and troubleshooting.



Figure 9. Cisco Nexus 93108TC-FX3P

The **Out-of-Band (Management) Network** encompasses all AI Compute node and BlueField-3 DPU base management controllers (BMC) as well as the network infrastructure management.

The following table lists the configuration of the Cisco Nexus 93108TC-FX3P.

Table 4. Cisco Nexus 93108TC-FX3P configuration

Part Number	Description	Quantity
7DL8CTO1WW	Cisco Nexus 9300-FX3 Series Switch (N9K-C93108TC-FX3)	
C5PB	Cisco Nexus 9300-GX2 Series Switch (N9K-C9364D-GX2A)	2
C6FK	Mode selection between ACI and NXOS (MODE-NXOS)	2
6252	2.5m, 16A/100-250V, C19 to C20 Jumper Cord	4
C1P1TN9300XF-5Y	5 Years (60 months) Cisco software Premier license	2

Cisco Nexus Dashboard

Cisco Nexus Dashboard, included with every Cisco Nexus 9000 switch tiered licensing purchase, serves as a centralized hub that unifies disparate network configurations and views from multiple switches and data centers. For AI/ML fabric operations, it acts as the ultimate command center, from the initial setup of AI/ML fabric automation to continuous fabric analytics within few clicks.

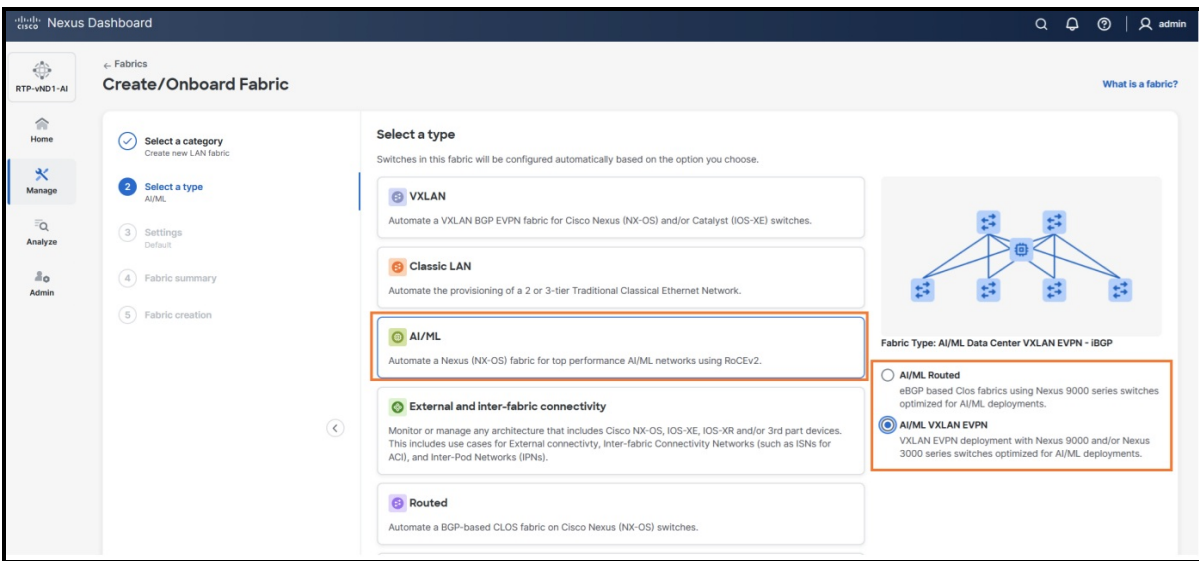


Figure 10. AI/ML Fabric Workflow on Nexus Dashboard

Key capabilities such as congestion scoring, PFC/ECN statistics, and microburst detection empower organizations to proactively identify and address performance bottlenecks for their AI/ML backend infrastructure.

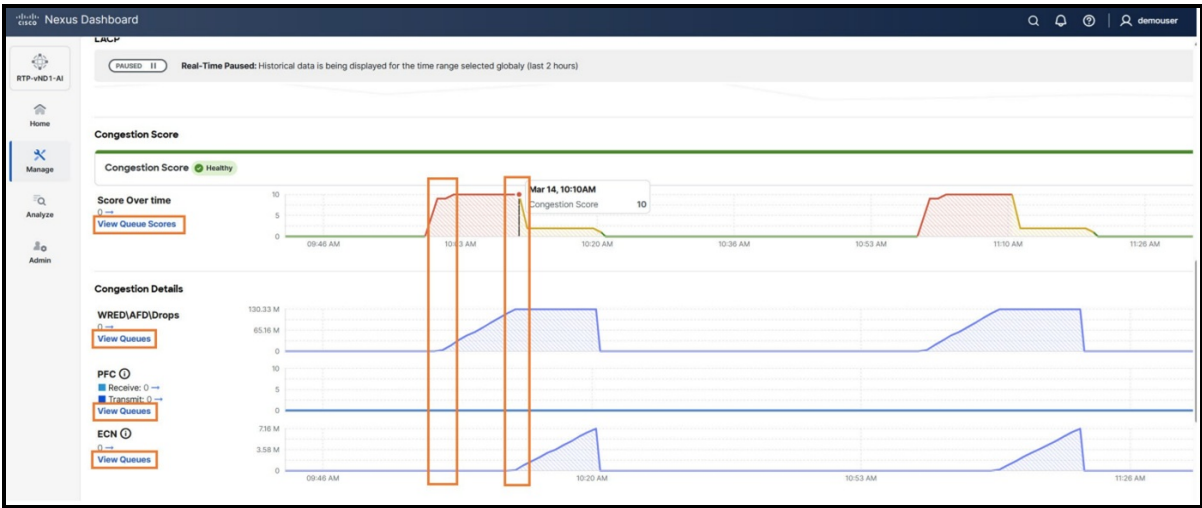


Figure 11. Congestion Score and Congestion Details on Nexus Dashboard

Advanced features like anomaly detection, event correlation, and suggested remediation ensure networks are not only resilient but also self-healing, minimizing downtime and accelerating issue resolution.

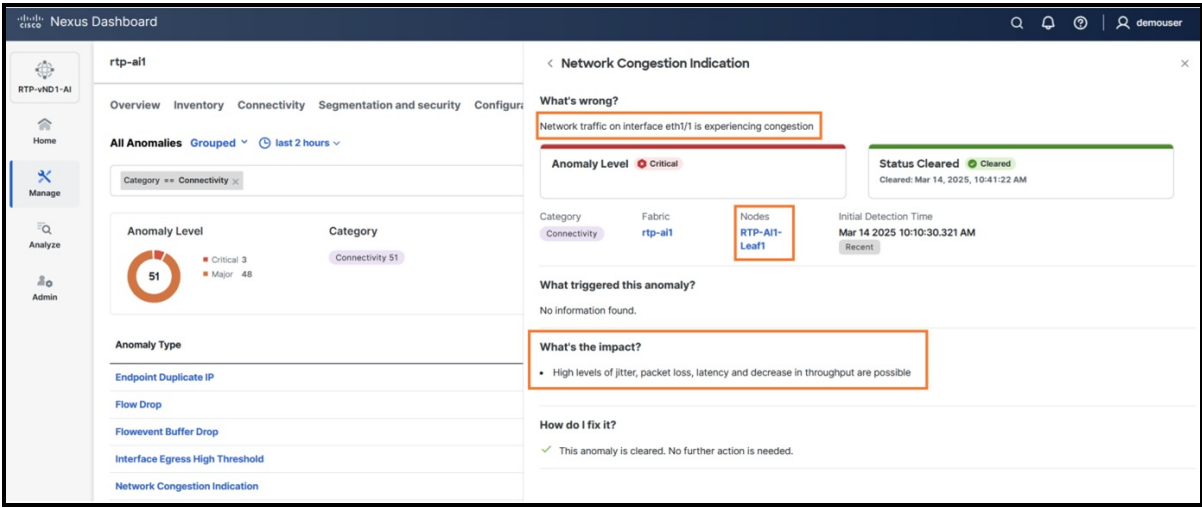


Figure 12. Anomaly Detection on Nexus Dashboard

Purpose-built to handle the high demands of AI workloads, the Cisco Nexus Dashboard transforms network management into a seamless, data-driven experience, unlocking the full potential of AI/ML fabrics.

Lenovo EveryScale Solution

The Server and Networking components and Operating System can come together as a [Lenovo EveryScale Solution](#). It is a framework for designing, manufacturing, integrating and delivering data center solutions, with a focus on High Performance Computing (HPC), Technical Computing, and Artificial Intelligence (AI) environments.

Lenovo EveryScale provides [Best Recipe](#) guides to warrant [interoperability](#) of hardware, software and firmware among a variety of Lenovo and third-party components.

Addressing specific needs in the data center, while also optimizing the solution design for application performance, requires a significant level of effort and expertise. Customers need to choose the right hardware and software components, solve interoperability challenges across multiple vendors, and determine optimal firmware levels across the entire solution to ensure operational excellence, maximize performance, and drive best total cost of ownership.

Lenovo EveryScale reduces this burden on the customer by pre-testing and validating a large selection of Lenovo and third-party components, to create a “Best Recipe” of components and firmware levels that work seamlessly together as a solution. From this testing, customers can be confident that such a best practice solution will run optimally for their workloads, tailored to the client’s needs.

In addition to interoperability testing, Lenovo EveryScale hardware is pre-integrated, pre-cabled, pre-loaded with the best recipe and optionally an OS-image and tested at the rack level in manufacturing, to ensure a reliable delivery and minimize installation time in the customer data center.

Scalability

A fundamental principle of the solution design philosophy is its ability to support any scale necessary to achieve a particular objective.

In a typical Enterprise AI deployment initially the AI environment is being used with a single use case, like for example an Enterprise RAG pipeline which can connect a Large Language Model (LLM) to Enterprise data for actionable insights grounded in relevant data.

In its simplest form, leveraging the NVIDIA Blueprint for Enterprise RAG pipeline involves three NVIDIA Inference Microservices: a Retriever, a Reranker, and the actual LLM. This setup requires a minimum of three GPUs.

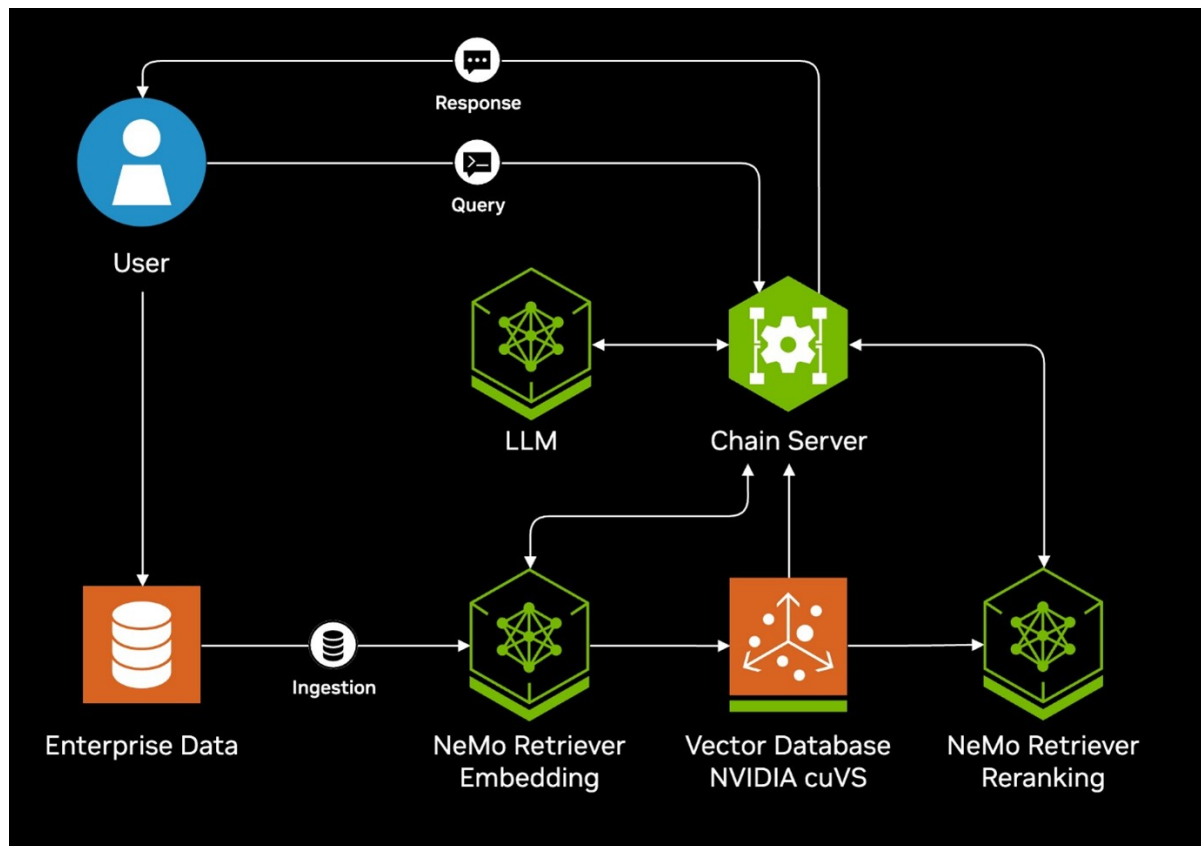


Figure 13. NVIDIA Blueprint Architecture Diagram

As the deployment of AI within the company continues to grow, the AI environment will be adapted to incorporate additional use cases, including Assistants or AI Agents. Additionally, it has the capacity to scale to support an increasing number of Microservices. Ultimately, most companies will maintain multiple AI environments operating simultaneously with their AI Agents working in unison.

The Lenovo Hybrid AI 285 with Cisco Networking platform has been designed to meet the customer where they are at with their AI application and then seamlessly scale with them through their AI integration. This is achieved through the introduction of the following:

- [Entry and AI Starter Kit Deployments](#)
- [Scalable Unit Deployment](#)
- [Custom Deployment](#)

A visual representation of the sizing and scaling of the platform is shown in the figure below.

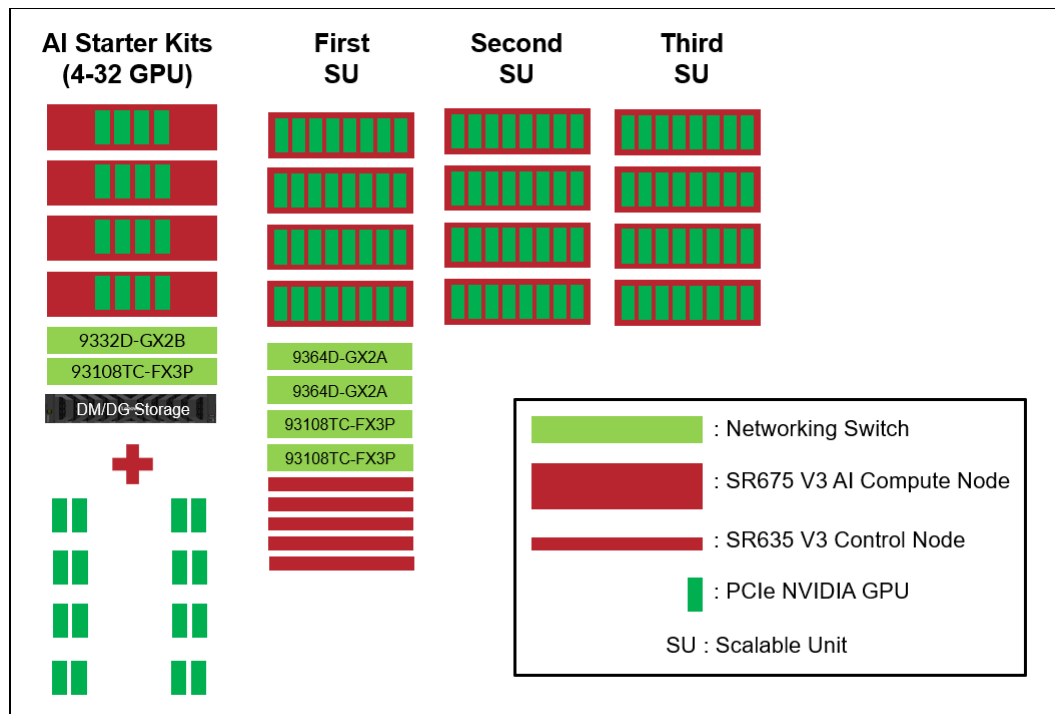


Figure 14. Lenovo Hybrid AI 285 with Cisco Networking Scaling

Entry and AI Starter Kit Deployments

Entry deployment sizings are for customers that want to deploy their initial AI factory with 4-16x GPUs. Entry deployments have one or two SR675 V3 servers, with 8x GPUs per server (AI Compute Nodes). With two servers configured, the two servers are connected directly via the installed NVIDIA ConnectX-7 or NVIDIA BlueField adapters.

If additional networking is required or additional storage is required, then use the AI Starter Kit deployment, which supports up to 4x servers and up to 32x GPUs. The AI Starter Kit uses NVIDIA networking switches and ThinkSystem DM or DG external storage.

The following sections describe these deployments:

- [Entry sizing](#)
- [AI Starter Kits](#)

Entry sizing

Entry sizing starts with a single AI Compute node, equipped with four GPUs. Such entry deployments are ideal for development, application trials, or small-scale use, reducing hardware costs, control plane overhead, and networking complexity. With all components on one node, management and maintenance are simplified.

Entry sizing can also support two AI Compute nodes, directly connected together and without the need for external networking switches, to scale up to 16 GPUs if fully populated (2 nodes, 8x GPUs per node). The two nodes connect to the rest of your data center using existing networking in your data center.

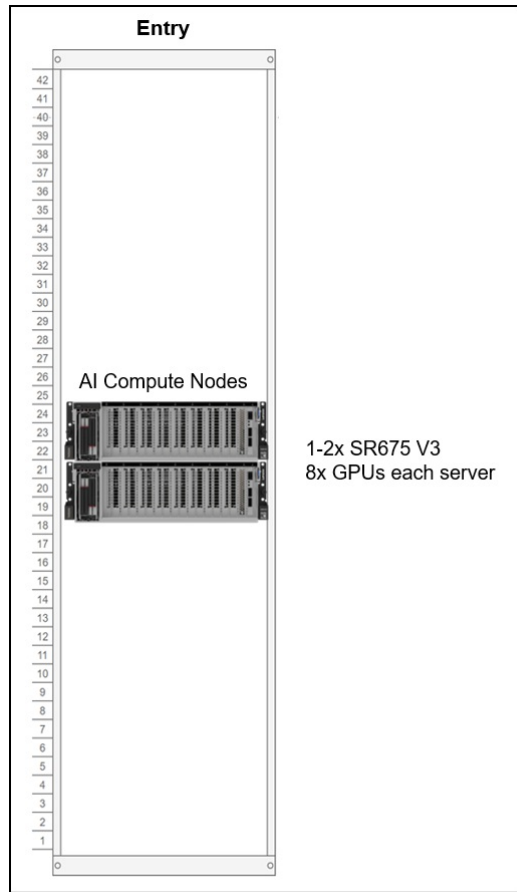


Figure 15. Entry Deployment Rack View

Table 5. Entry sizing

	4-8x GPUs	8-16x GPUs
Compute	1x SR675 V3	2x SR675 V3
Network adapters per server	Minimum ratio of 1 CX7 per 2 GPUs	Minimum ratio of 1 CX7 per 2 GPUs

AI Starter Kits

For customers who want storage and/or networking in the Entry sizing, Lenovo and Cisco worked to develop AI Starter Kits with Cisco networking which allows up to 32 GPUs across 4 nodes, slightly more than the base 285 starter kit sizing which allows for up to only 24 GPUs. This sizing is for customers who do not plan to scale above 32 GPUs in the near future but still need an end-to-end solution for compute and storage.

Networking between the nodes is implemented using the Cisco 9332D-GX2B 200GbE switches and NVIDIA ConnectX-7 dual-port 200Gb adapters in each server.

Storage is implemented using either ThinkSystem DM or ThinkSystem DG Storage Arrays. Features include:

- Easy to deploy and scale for performance or capacity
- Unified file, object, and block eliminates AI data silos
- High performance NVMe flash and GPUDirect enable faster time to insights
- Confidently use production data to fine tune models with advanced data management features

The table and figure below show the hardware involved in various sizes of AI Starter Kit deployments.

Table 6. AI Starter Kit sizing with Cisco Networking

	4-8x GPUs	8-16x GPUs	32x GPUs
Compute	1x SR675 V3	2x SR675 V3	4x SR675 V3
Storage	DG5200	DM7200F	DM7200F
Network adapters per server	5x CX-7	5x CX-7	5x CX-7
Networking	9332D-GX2B 93108TC-FX3P	9332D-GX2B 93108TC-FX3P	9332D-GX2B 93108TC-FX3P

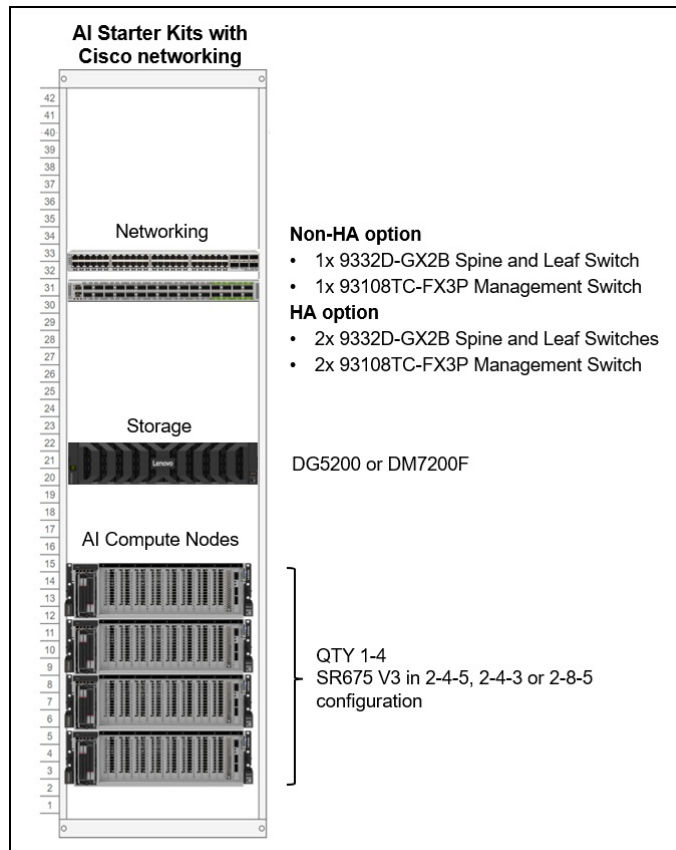


Figure 16. AI Starter Kits

Scalable Unit Deployment

For configurations beyond two nodes, it is advisable to deploy a full Scalable Unit along with the necessary network and service infrastructure, providing a foundation for further growth in enterprise use cases.

The first SU consists of up to four AI Compute nodes, minimum five service nodes, and networking switches. When additional AI Compute Nodes are required, additional SUs of four AI Compute Nodes can be added.

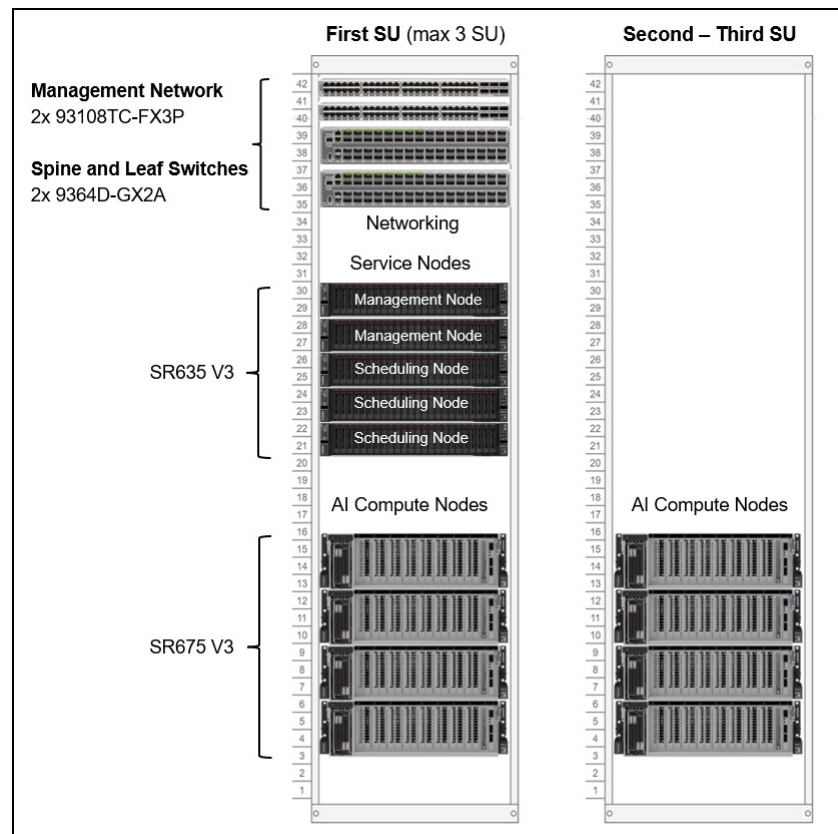


Figure 17. Scalable Unit Deployment

Networking is implemented using Cisco 93108TC-FX3P switches and CX-7 adapters in the AI Compute Nodes. The combination of these two pieces allows the user to take advantage of NVIDIA's Spectrum-X networking, an Ethernet platform that delivers the highest performance for AI, machine learning, and natural language processing.

The networking decision depends on whether the platform is designed to support up to three Scalable Units in total, and whether it will handle exclusively inference workloads or also encompass future fine-tuning and re-training activities. Subsequently, the solution can be expanded seamlessly without downtime by incorporating additional Scalable Units, ultimately reaching a total of three as needed.

Custom Deployment

For high-end scenarios requiring more than eight scalable units, the network can be custom designed to any required size. Lenovo will develop a fully bespoke solution tailored to match the workflow and workload requirements in that case.

Performance

This document is meant to be used in tandem with the [Hybrid AI 285 platform guide](#) please see this section there.

AI Software Stack

This document is meant to be used in tandem with the [Hybrid AI 285 platform guide](#) please see this section there.

Storage Connectivity

This document is meant to be used in tandem with the [Hybrid AI 285 platform guide](#) please see this section there.

Lenovo AI Center of Excellence

This document is meant to be used in tandem with the [Hybrid AI 285 platform guide](#) please see this section there.

AI Services

This document is meant to be used in tandem with the [Hybrid AI 285 platform guide](#) please see this section there.

Lenovo TruScale

This document is meant to be used in tandem with the [Hybrid AI 285 platform guide](#) please see this section there.

Lenovo Financial Services

This document is meant to be used in tandem with the [Hybrid AI 285 platform guide](#) please see this section there.

Bill of Materials - 3 Scalable Unit (3 SU)

This section provides an example Bill of Materials (BoM) of one Scaleable Unit (SU) deployment with NVIDIA Spectrum-X.

This example BoM includes:

- 12x Lenovo ThinkSystem SR675 V3 with 8x NVIDIA H200 NVL GPUs per server (4 Servers/Scalable Unit)
- 5x Lenovo ThinkSystem SR635 V3
- 2x Cisco 9364D-GX2A Switches
- 2x Cisco 93108TC-FX3P Switches

Storage is optional and not included in this 3 SU BoM.

In this section:

- [3SU: ThinkSystem SR675 V3 BoM](#)
- [3SU: ThinkSystem SR635 V3 BoM](#)
- [3SU: Cisco 9364D-GX2A Switch BoM](#)
- [3SU: Cisco 93108TC-FX3P Switch BoM](#)
- [3SU: Power Distribution Unit \(PDU\) BoM](#)
- [3SU: Rack Cabinet BoM](#)
- [3SU: XClarity Software BoM](#)
- [3SU: Cables and Transceivers BoM](#)

3SU: ThinkSystem SR675 V3 BoM

Table 7. ThinkSystem SR675 V3 BoM

Part Number	Product Description	Qty per System	Total Qty
7D9RCTOLWW	ThinkSystem SR675 V3		12
BR7F	ThinkSystem SR675 V3 8DW PCIe GPU Base	1	12
C3EF	ThinkSystem SR675 V3 System Board v2	1	12
C2AL	ThinkSystem AMD EPYC 9535 64C 300W 2.4GHz Processor	2	24
C0CK	ThinkSystem 64GB TruDDR5 6400MHz (2Rx4) RDIMM-A	24	288
BR7S	ThinkSystem SR675 V3 Switched 4x16 PCIe DW GPU Direct RDMA Riser	2	24
C3V3	ThinkSystem NVIDIA H200 NVL 141GB PCIe GPU Gen5 Passive GPU	8	96
C3V0	ThinkSystem NVIDIA 4-way bridge for H200 NVL	2	24
BR7H	ThinkSystem SR675 V3 2x16 PCIe Front IO Riser	1	12
C2RK	ThinkSystem SR675 V3 2 x16 Switch Cabled PCIe Rear IO Riser	2	24
BQBN	ThinkSystem NVIDIA ConnectX-7 NDR200/200GbE QSFP112 2-port PCIe Gen5 x16 Adapter	5	60
BM8X	ThinkSystem M.2 SATA/x4 NVMe 2-Bay Adapter	1	12
BT7P	ThinkSystem Raid 540-8i for M.2/7MM NVMe boot Enablement	1	12
BXMH	ThinkSystem M.2 PM9A3 960GB Read Intensive NVMe PCIe 4.0 x4 NHS SSD	2	24
BTMB	ThinkSystem 1x4 E3.S Backplane	1	12
C1AB	ThinkSystem E3.S PM9D3a 3.84TB Read Intensive NVMe PCIe 5.0 x4 HS SSD	2	24
BK1E	ThinkSystem SR670 V2/ SR675 V3 OCP Enablement Kit	1	12
C5WW	ThinkSystem SR675 V3 Dual Rotor System High Performance Fan	5	60
BFD6	ThinkSystem SR670 V2/ SR675 V3 Power Mezzanine Board	1	12
BE0D	N+1 Redundancy With Over-Subscription	1	12
BKTJ	ThinkSystem 2600W 230V Titanium Hot-Swap Gen2 Power Supply	4	48
6252	2.5m, 16A/100-250V, C19 to C20 Jumper Cord	4	48
C3KA	ThinkSystem SR670 V2/SR675 V3 Heavy Systems Toolless Slide Rail Kit	1	12
BFNU	ThinkSystem SR670 V2/ SR675 V3 Intrusion Cable	1	12
BR7U	ThinkSystem SR675 V3 Root of Trust Module	1	12
BFTH	ThinkSystem SR670 V2/ SR675 V3 Front Operator Panel ASM	1	12
5PS7B09631	5Yr Premier NBD Resp + KYD SR675 V3	1	12

3SU: ThinkSystem SR635 V3 BoM

Table 8. ThinkSystem SR635 V3 BoM

Part Number	Product Description	Qty per System	Total Qty
7D9GCTO1WW	Server : ThinkSystem SR635 V3 - 3yr Warranty		5
BLK4	ThinkSystem V3 1U 10x2.5" Chassis	1	5
BVGL	Data Center Environment 30 Degree Celsius / 86 Degree Fahrenheit	1	5
C2AQ	ThinkSystem AMD EPYC 9335 32C 210W 3.0GHz Processor	1	5
BQ26	ThinkSystem SR645 V3/SR635 V3 1U High Performance Heatsink	1	5
C1PL	ThinkSystem 32GB TruDDR5 6400MHz (1Rx4) RDIMM-A	12	60
BC4V	Non RAID NVMe	1	5
C0ZU	ThinkSystem 2.5" U.2 VA 3.84TB Read Intensive NVMe PCIe 5.0 x4 HS SSD	2	10
BPC9	ThinkSystem 1U 4x 2.5" NVMe Gen 4 Backplane	1	5
B5XJ	ThinkSystem M.2 SATA/NVMe 2-Bay Adapter	1	5
BTTY	M.2 NVMe	1	5
BKSR	ThinkSystem M.2 7450 PRO 960GB Read Intensive NVMe PCIe 4.0 x4 NHS SSD	2	10
BQBN	ThinkSystem NVIDIA ConnectX-7 NDR200/200GbE QSFP112 2-port PCIe Gen5 x16 Adapter	1	5
BLK7	ThinkSystem SR635 V3/SR645 V3 x16 PCIe Gen5 Riser 1	1	5
BLK9	ThinkSystem V3 1U MS LP+LP BF Riser Cage	1	5
BNFG	ThinkSystem 750W 230V/115V Platinum Hot-Swap Gen2 Power Supply v3	2	10
BH9M	ThinkSystem V3 1U Performance Fan Option Kit v2	7	35
BLKD	ThinkSystem 1U V3 10x2.5" Media Bay w/ Ext. Diagnostics Port	1	5
7Q01CTS2WW	5Yr Premier NBD Resp + KYD SR635 V3	1	5

3SU: Cisco 9364D-GX2A Switch BoM

Table 9. Cisco 9364D-GX2A Switch BoM

Part Number	Product Description	Total Qty
7DLKCTO1WW	Cisco Nexus 9300-GX2 Series Switch (N9K-C9364D-GX2A)	
C5P0	Cisco Nexus 9300-GX2 Series Switch (N9K-C9364D-GX2A)	2
C6FK	Mode selection between ACI and NXOS (MODE-NXOS)	2
6252	2.5m, 16A/100-250V, C19 to C20 Jumper Cord	4
C1P1TN9300XF2-5Y	5 Years (60 months) Cisco software Premier license	2

3SU: Cisco 93108TC-FX3P Switch BoM

Table 10. Cisco 93108TC-FX3P Switch BoM

Part Number	Description	Total Qty
7DL8CTO1WW	Cisco Nexus 9300-FX3 Series Switch (N9K-C93108TC-FX3)	
C5PB	Cisco Nexus 9300-GX2 Series Switch (N9K-C9364D-GX2A)	2
C6FK	Mode selection between ACI and NXOS (MODE-NXOS)	2
6252	2.5m, 16A/100-250V, C19 to C20 Jumper Cord	4
C1P1TN9300XF-5Y	5 Years (60 months) Cisco software Premier license	2

3SU: Power Distribution Unit (PDU) BoM

Table 11. Power Distribution Unit (PDU) BoM

Part Number	Product Description	Qty per System	Total Qty
7DGMCTO1WW	-SB- 0U 18 C13/C15 and 18 C13/C15/C19 Switched and Monitored 63A 3 Phase WYE PDU v2		2

3SU: Rack Cabinet BoM

Table 12. Rack Cabinet BoM

Part Number	Product Description	Qty per System	Total Qty
1410O42	Lenovo EveryScale 42U Onyx Heavy Duty Rack Cabinet		1
BHC4	Lenovo EveryScale 42U Onyx Heavy Duty Rack Cabinet	1	1
BJPD	21U Front Cable Management Bracket	2	2
BHC7	ThinkSystem 42U Onyx Heavy Duty Rack Side Panel	2	2
BJPA	ThinkSystem 42U Onyx Heavy Duty Rack Rear Door	1	1
5AS7B07693	Lenovo EveryScale Rack Setup Services	1	1

3SU: XClarity Software BoM

Table 13. XClarity Software BoM

Part Number	Product Description	Qty per System	Total Qty
SBCV	Lenovo XClarity XCC2 Platinum Upgrade (FOD)		3
00MT203	Lenovo XClarity Pro, Per Managed Endpoint w/5 Yr SW S&S		5

3SU: Cables and Transceivers BoM

Table 14. Cables and Transceivers BoM

Part Number	Product Description	Total Qty
P01DQ3007-07-R Example from luxshare-tech.com	400G QSFP-DD to 2 x 200G QSFP56 Active Optical Breakout Cable 7M	48
QSFP-200-CU3M From Cisco	200G QSFP56 to QSFP56, Passive Copper Cable 3m	10
QDD-400-CU2M From Cisco	400 Gbps, QSFP-DD to QSFPDD, DAC, 2M	10
7Z57A03562	Lenovo 3M Passive 100G QSFP28 DAC Cable	12

Bill of Materials – Starter Kit

This section provides an example Bill of Materials (BoM) of Starter Kit deployment with Cisco switches.

This example BoM includes:

- 4x Lenovo ThinkSystem SR675 V3 with 8 × NVIDIA H200 NVL GPUs per server (The servers can be configured with less GPUs for an underpopulated configuration)
- 2x Cisco 9364D-GX2A Switches
- 2x Cisco 93108TC-FX3P Switches
- 1x Lenovo ThinkSystem DM7200 Storage

In this section:

- [Starter: ThinkSystem SR675 V3 BoM](#)
- [Starter: ThinkSystem DM7200F Storage BoM](#)
- [Starter: Cisco 9364D-GX2A Switch BoM](#)
- [Starter: Cisco 93108TC-FX3P Switch BoM](#)
- [Starter: Cables and Transceivers BoM](#)

Starter: ThinkSystem SR675 V3 BoM

Table 15. ThinkSystem SR675 V3 BoM

Part Number	Product Description	Qty per System	Total Qty
7D9RCTOLWW	ThinkSystem SR675 V3		4
BR7F	ThinkSystem SR675 V3 8DW PCIe GPU Base	1	4
C3EF	ThinkSystem SR675 V3 System Board v2	1	4
C2AL	ThinkSystem AMD EPYC 9535 64C 300W 2.4GHz Processor	2	8
C0CK	ThinkSystem 64GB TruDDR5 6400MHz (2Rx4) RDIMM-A	24	96
BR7S	ThinkSystem SR675 V3 Switched 4x16 PCIe DW GPU Direct RDMA Riser	2	8
C3V3	ThinkSystem NVIDIA H200 NVL 141GB PCIe GPU Gen5 Passive GPU	8	32
C3V0	ThinkSystem NVIDIA 4-way bridge for H200 NVL	2	8
BR7H	ThinkSystem SR675 V3 2x16 PCIe Front IO Riser	1	4
C2RK	ThinkSystem SR675 V3 2 x16 Switch Cabled PCIe Rear IO Riser	2	4
BQBN	ThinkSystem NVIDIA ConnectX-7 NDR200/200GbE QSFP112 2-port PCIe Gen5 x16 Adapter	5	8
BM8X	ThinkSystem M.2 SATA/x4 NVMe 2-Bay Adapter	1	16
BT7P	ThinkSystem Raid 540-8i for M.2/7MM NVMe boot Enablement	1	4
BXMH	ThinkSystem M.2 PM9A3 960GB Read Intensive NVMe PCIe 4.0 x4 NHS SSD	2	4
BTMB	ThinkSystem 1x4 E3.S Backplane	1	8
C1AB	ThinkSystem E3.S PM9D3a 3.84TB Read Intensive NVMe PCIe 5.0 x4 HS SSD	2	4
BK1E	ThinkSystem SR670 V2/ SR675 V3 OCP Enablement Kit	1	8
C5WW	ThinkSystem SR675 V3 Dual Rotor System High Performance Fan	5	4
BFD6	ThinkSystem SR670 V2/ SR675 V3 Power Mezzanine Board	1	20
BE0D	N+1 Redundancy With Over-Subscription	1	4
BKTJ	ThinkSystem 2600W 230V Titanium Hot-Swap Gen2 Power Supply	4	4
6252	2.5m, 16A/100-250V, C19 to C20 Jumper Cord	4	16
C3KA	ThinkSystem SR670 V2/SR675 V3 Heavy Systems Toolless Slide Rail Kit	1	16
BFNU	ThinkSystem SR670 V2/ SR675 V3 Intrusion Cable	1	4
BR7U	ThinkSystem SR675 V3 Root of Trust Module	1	4
BFTH	ThinkSystem SR670 V2/ SR675 V3 Front Operator Panel ASM	1	4
5PS7B09631	5Yr Premier NBD Resp + KYD SR675 V3	1	4

Starter: ThinkSystem DM7200F Storage BoM

Table 16. ThinkSystem DM7200F Storage BoM

Part Number	Product Description	Total Qty
Part Number	Product Description	Qty
7DJ3CTO1WW	Controller : Lenovo ThinkSystem DM7200F All Flash Array	1
BF3C	Lenovo ThinkSystem Storage 2U NVMe Chassis	1
BWU8	Storage Complete Bundle Offering	1
C4A4	Lenovo ThinkSystem DM7200 Series Controller, 128GB	2
C3XK	Lenovo ThinkSystem 30.7TB (2x 15.36TB NVMe SED) Drive Pack	9
C4AA	Lenovo ThinkSystem Storage 100Gb 2 port Ethernet, RoCE Adapter (Host/Cluster)	2
C4AA	Lenovo ThinkSystem Storage 100Gb 2 port Ethernet, RoCE Adapter (Host/Cluster)	4
C4AG	Lenovo ThinkSystem Storage ONTAP 9.16 Software Encryption - IPAv2	1
B0W1	3 Years	1
C6S2	Premier 24x7 4hr Response and KYD	1
C48T	Configured with Lenovo ThinkSystem DM7200F 3Yr Warranty	1
BWUE	Storage Encryption Bundle License Key - RoW	2
BWUC	Storage Complete Bundle License Key	2
C49B	Lenovo ThinkSystem DM/DG Series Jupiter All Flash Ship Kit - Multi-Language	1
B6Y6	Lenovo ThinkSystem NVMe Rail Kit 4 post	1
7S0SCTOMWW	ThnkSys DM7200F 7DJ3 SWLicense	1
SDJE	Lenovo ThinkSystem DM7200F NVMe SSD Unified Complete SW License with 3 Years Support, Per 0.1TB	2765
5641PX3	XClarity Pro, Per Endpoint w/3 Yr SW S&S	1
1340	Lenovo XClarity Pro, Per Managed Endpoint w/3 Yr SW S&S	1
3444	Registration only	1
5WS7C06619	3Yr Premier 24x7 4Hr Resp DM7200F+KYD	1
5WS7C07259	3Yr Premier 24x7 4Hr Resp+KYD (0.1TB NVMe TLC)	2765

Starter: Cisco 9364D-GX2A Switch BoM

Table 17. Cisco 9364D-GX2A Switch BoM

Part Number	Product Description	Total Qty
7DLKCTO1WW	Cisco Nexus 9300-GX2 Series Switch (N9K-C9364D-GX2A)	
C5P0	Cisco Nexus 9300-GX2 Series Switch (N9K-C9364D-GX2A)	2
C6FK	Mode selection between ACI and NXOS (MODE-NXOS)	2
6252	2.5m, 16A/100-250V, C19 to C20 Jumper Cord	4
C1P1TN9300XF2-5Y	5 Years (60 months) Cisco software Premier license	2

Starter: Cisco 93108TC-FX3P Switch BoM

Table 18. Cisco 93108TC-FX3P Switch BoM

Part Number	Description	Total Qty
7DL8CTO1WW	Cisco Nexus 9300-FX3 Series Switch (N9K-C93108TC-FX3)	
C5PB	Cisco Nexus 9300-GX2 Series Switch (N9K-C9364D-GX2A)	2
C6FK	Mode selection between ACI and NXOS (MODE-NXOS)	2
6252	2.5m, 16A/100-250V, C19 to C20 Jumper Cord	4
C1P1TN9300XF-5Y	5 Years (60 months) Cisco software Premier license	2

Starter: Cables and Transceivers BoM

Table 19. Cables and Transceivers BoM

Part Number	Product Description	Total Qty
QSFP-200-CU3M From Cisco	200G QSFP56 to QSFP56, Passive Copper Cable 3m	20
QDD-400-CU2M From Cisco	400 Gbps, QSFP-DD to QSFPDD, DAC, 2M	10
7Z57A03562	Lenovo 3M Passive 100G QSFP28 DAC Cable	12

Seller training courses

The following sales training courses are offered for employees and partners (login required). Courses are listed in date order.

1. VTT AI: Introducing the Lenovo Hybrid AI 285 Platform April 2025

2025-04-30 | 60 minutes | Employees Only

The Lenovo Hybrid AI 285 Platform enables enterprises of all sizes to quickly deploy AI infrastructures supporting use cases as either new greenfield environments or as an extension to current infrastructures. The 285 Platform enables the use of the NVIDIA AI Enterprise software stack. The AI Hybrid 285 platform is the perfect foundation supporting Lenovo Validated Designs.

- Technical overview of the Hybrid AI 285 platform
- AI Hybrid platforms as infrastructure frameworks for LVDs addressing data center-based AI solutions.
- Accelerate AI adoption and reduce deployment risks

Tags: Artificial Intelligence (AI), Nvidia, Technical Sales, Lenovo Hybrid AI 285

Published: 2025-04-30

Length: 60 minutes

Start the training:

Employee link: Grow@Lenovo

Course code: DVAI215

Related publications and links

For more information, see these resources:

- Lenovo EveryScale support page:
<https://datacentersupport.lenovo.com/us/en/solutions/ht505184>
- x-config configurator:
<https://lesc.lenovo.com/products/hardware/configurator/worldwide/bhui/asit/x-config.jnlp>
- Implementing AI Workloads using NVIDIA GPUs on ThinkSystem Servers:
<https://lenovopress.lenovo.com/lp1928-implementing-ai-workloads-using-nvidia-gpus-on-thinksystem-servers>
- Making LLMs Work for Enterprise Part 3: GPT Fine-Tuning for RAG:
<https://lenovopress.lenovo.com/lp1955-making-llms-work-for-enterprise-part-3-gpt-fine-tuning-for-rag>
- Lenovo to Deliver Enterprise AI Compute for NetApp AI Pod Through Collaboration with NetApp and NVIDIA
<https://lenovopress.lenovo.com/lp1962-lenovo-to-deliver-enterprise-ai-compute-for-netapp-ai-pod-nvidia>

Related product families

Product families related to this document are the following:

- [AI Servers](#)
- [Artificial Intelligence](#)
- [ThinkSystem SR635 V3 Server](#)
- [ThinkSystem SR675 V3 Server](#)

Notices

Lenovo may not offer the products, services, or features discussed in this document in all countries. Consult your local Lenovo representative for information on the products and services currently available in your area. Any reference to a Lenovo product, program, or service is not intended to state or imply that only that Lenovo product, program, or service may be used. Any functionally equivalent product, program, or service that does not infringe any Lenovo intellectual property right may be used instead. However, it is the user's responsibility to evaluate and verify the operation of any other product, program, or service. Lenovo may have patents or pending patent applications covering subject matter described in this document. The furnishing of this document does not give you any license to these patents. You can send license inquiries, in writing, to:

Lenovo (United States), Inc.
8001 Development Drive
Morrisville, NC 27560
U.S.A.
Attention: Lenovo Director of Licensing

LENOVO PROVIDES THIS PUBLICATION "AS IS" WITHOUT WARRANTY OF ANY KIND, EITHER EXPRESS OR IMPLIED, INCLUDING, BUT NOT LIMITED TO, THE IMPLIED WARRANTIES OF NON-INFRINGEMENT, MERCHANTABILITY OR FITNESS FOR A PARTICULAR PURPOSE. Some jurisdictions do not allow disclaimer of express or implied warranties in certain transactions, therefore, this statement may not apply to you.

This information could include technical inaccuracies or typographical errors. Changes are periodically made to the information herein; these changes will be incorporated in new editions of the publication. Lenovo may make improvements and/or changes in the product(s) and/or the program(s) described in this publication at any time without notice.

The products described in this document are not intended for use in implantation or other life support applications where malfunction may result in injury or death to persons. The information contained in this document does not affect or change Lenovo product specifications or warranties. Nothing in this document shall operate as an express or implied license or indemnity under the intellectual property rights of Lenovo or third parties. All information contained in this document was obtained in specific environments and is presented as an illustration. The result obtained in other operating environments may vary. Lenovo may use or distribute any of the information you supply in any way it believes appropriate without incurring any obligation to you.

Any references in this publication to non-Lenovo Web sites are provided for convenience only and do not in any manner serve as an endorsement of those Web sites. The materials at those Web sites are not part of the materials for this Lenovo product, and use of those Web sites is at your own risk. Any performance data contained herein was determined in a controlled environment. Therefore, the result obtained in other operating environments may vary significantly. Some measurements may have been made on development-level systems and there is no guarantee that these measurements will be the same on generally available systems. Furthermore, some measurements may have been estimated through extrapolation. Actual results may vary. Users of this document should verify the applicable data for their specific environment.

© Copyright Lenovo 2025. All rights reserved.

This document, LP2236, was created or updated on June 17, 2025.

Send us your comments in one of the following ways:

- Use the online Contact us review form found at:
<https://lenovopress.lenovo.com/LP2236>
- Send your comments in an e-mail to:
comments@lenovopress.com

This document is available online at <https://lenovopress.lenovo.com/LP2236>.

Trademarks

Lenovo and the Lenovo logo are trademarks or registered trademarks of Lenovo in the United States, other countries, or both. A current list of Lenovo trademarks is available on the Web at <https://www.lenovo.com/us/en/legal/copytrade/>.

The following terms are trademarks of Lenovo in the United States, other countries, or both:

Lenovo®

Lenovo Hybrid AI Advantage

ThinkAgile®

ThinkSystem®

XClarity®

The following terms are trademarks of other companies:

AMD and AMD EPYC™ are trademarks of Advanced Micro Devices, Inc.

Intel® is a trademark of Intel Corporation or its subsidiaries.

Linux® is the trademark of Linus Torvalds in the U.S. and other countries.

IBM® is a trademark of IBM in the United States, other countries, or both.

Other company, product, or service names may be trademarks or service marks of others.