



Optimizing AI Storage Fabrics

NVIDIA Spectrum-X Accelerates AI Storage Networks

White Paper

Table of Contents

- Abstract..... 1
- AI Clusters Need Fast Storage 2
 - Impact of Slow AI Storage..... 2
 - Challenges in AI Storage Networks 3
 - Optimized Network Support for RDMA 4
 - Typical AI Storage Fabric for Training..... 5
- NVIDIA Spectrum-X Technology 7
 - NVIDIA Spectrum-X SuperNIC 8
 - NVIDIA SN5600 Ethernet Switch..... 9
- Accelerating AI Storage Networks with NVIDIA Spectrum-X 10
- AI Storage Fabric Benchmark Results..... 12
 - Spectrum-X Performance Results from Israel-1 12
 - Spectrum-X Improves AI Storage Visibility..... 16
 - Spectrum-X Improves AI Storage Fabric Resiliency..... 17
 - Virtualizing AI Storage Fabrics with Digital Twins in NVIDIA AIR..... 18
 - Benefits of Virtualizing the AI Storage System..... 18
- Summary20

List of Figures

Figure 1. Example AI Storage Fabric Based on InfiniBand.....	6
Figure 2. NVIDIA Spectrum-X BlueField-3 Ethernet SuperNIC	9
Figure 3. NVIDIA SN5600 Ethernet Switch.....	9
Figure 4. Israel-1 AI Supercomputer: Validated / Optimized / Quantified	12
Figure 5. Israel-1 SU E/W Topology.....	13
Figure 6. Israel-1 Storage Fabric Topology	14
Figure 7. Spectrum-X Increases Storage Bandwidth Through RoCE Adaptive Routing.	15
Figure 8. NetQ GUI Reporting.....	17

List of Tables

Table 1. Spectrum-X Components	7
--------------------------------------	---

Abstract

As AI models grow increasingly complex and data-intensive, the demand for high-performance training infrastructure has never been greater. Traditional storage networks often struggle to keep pace with the massive data throughput and low latency requirements of modern AI workloads, leading to bottlenecks that hinder training speed and model efficiency. This white paper explores the critical role of accelerated storage networks in optimizing AI training clusters, specifically highlighting the capabilities of NVIDIA® Spectrum™-X.

The NVIDIA Spectrum-X platform provides an advanced solution for enhancing storage performance by delivering high-speed Ethernet switching, a scalable architecture, and low-latency data access. A key feature of Spectrum-X is its adaptive routing technology, which improves storage read and write I/O performance by up to 1.6x, effectively eliminating ECMP collisions. Combined with a second key feature, Programmable Congestion Control (PCC), these technologies work together to prevent data bottlenecks and enhance overall training efficiency. Together, they also increase GPU utilization by ensuring faster data transfer and efficient checkpointing while significantly improving the resiliency of AI training clusters.

Spectrum-X seamlessly integrates into NVIDIA's full-stack AI solution, providing a cohesive and optimized environment for AI development and deployment. By aligning with NVIDIA's broader ecosystem of GPUs, NVIDIA NVLINK™, AI-optimized libraries, software, and frameworks, Spectrum-X ensures end-to-end performance optimization and simplified infrastructure management. Through case studies and performance analyses, this paper demonstrates the transformative impact of Spectrum-X on AI training, making it an indispensable tool for organizations looking to stay competitive in the AI landscape.

AI Clusters Need Fast Storage

In AI training clusters, storage plays a critical role in ensuring the efficient flow of data required for model training. AI training workloads often involve processing vast amounts of data, which must be rapidly accessed and manipulated to maximize the utilization of compute resources, particularly GPUs. Faster storage is essential for several reasons: it helps minimize GPU idle time, enabling continuous data prefetching and caching to keep the training pipeline fed with fresh data. This is crucial for training large language models (LLMs), managing extensive datasets, and handling batch loading, frequent checkpoints, and logging activities without delays.

Checkpointing becomes increasingly vital as cluster sizes and training times grow, increasing the likelihood of failures. Effective checkpointing ensures that training can resume from the last saved state, preventing significant loss of work and time. Additionally, fast storage is vital for supporting distributed training and parallelism, where multiple GPUs work simultaneously on different parts of a model or dataset. In these scenarios, slow storage creates bottlenecks that hinder performance and limit the scalability of AI workloads.

AI Inference use-cases rely heavily on fast storage IO. The speed of the storage per GPU can impact the duration of the ingestion phase in Retrieval Augmented Generative (RAG) AI workflows. Storage bandwidth also plays a critical role of the Key-Value (KV) cache, affecting both its efficiency and scalability.

Impact of Slow AI Storage

- > **Training Bottlenecks:** Slow storage often forces CPUs or GPUs to wait for data to be read from or written to storage, reducing overall training speed and hardware utilization. This leads to inefficient use of expensive compute resources.
- > **Longer Training Times:** Slow data loading and inefficient checkpoint saving significantly increase AI model training times. These delays accumulate over many iterations, resulting in prolonged training cycles that delay model deployment.
- > **Higher Costs:** Prolonged training times directly translate into higher computational costs, especially in cloud environments where costs are incurred based on compute

time. Inefficient storage can therefore lead to unnecessary expenditures and reduced cost-effectiveness of AI projects.

- > **Reduced Experimentation:** The increased time required for training makes it more challenging to experiment with different models or hyperparameters. This slows the pace of research and development, reducing the ability to innovate quickly and adapt models to new data or evolving requirements.

By addressing these storage performance challenges, AI training clusters can achieve higher performance, lower costs, and greater flexibility for experimentation, ultimately accelerating the development and deployment of AI models.

Challenges in AI Storage Networks

In an AI cluster, the architecture typically consists of GPU nodes and storage nodes interconnected through a leaf-spine network. This network architecture can be a dedicated fabric designed specifically for storage or a shared storage network that also manages traffic. The choice of network setup is critical to optimizing the efficiency and performance of AI training workloads. However, AI storage networks face several unique challenges that can significantly impact their performance and efficiency, especially when relying on traditional leaf-spine topologies.

1. **Inefficient Load Balancing for Large Data Flows:** Elephant flows—large, sustained data transfers between GPU and storage nodes—are a common feature of AI training clusters. Traditional leaf-spine networks typically use static Equal-Cost Multi-Path (ECMP) routing to balance traffic flow by flow across multiple paths. However, static ECMP is designed for many small, short-lived flows, lacks the ability to dynamically load balance traffic on a packet-by-packet basis, and cannot adapt to the varying demands of AI workloads. This leads to imbalanced load distribution, meaning some paths may become congested while others remain underutilized, resulting in decreased network performance and data transfer inefficiencies between GPU and storage nodes.
2. **Suboptimal Throughput Utilization:** The static nature of ECMP and its limitations in handling dynamic and large-scale data flows limits the throughput of standard leaf-spine networks. These networks often achieve only about 60% of the theoretical combined throughput of all leaf-to-spine connections—far below the performance required for AI workloads. High throughput is crucial for maximizing GPU utilization and minimizing training times. In scenarios where the leaf-spine network also handles management traffic on an in-band management network, competition for bandwidth further reduces the effective throughput available for storage operations.
3. **Inter-Switch Congestion and High Latency:** Uneven utilization of network paths due to static load balancing leads to inter-switch congestion, which not only reduces throughput but also increases latency. In an AI cluster, high latency disrupts the smooth flow of data required for efficient training, forcing GPUs to wait idly for data and slowing overall training speeds. This issue is further exacerbated in networks that combine storage and management traffic, where high-latency scenarios become more frequent.

4. **Frequent Link Issues in Dense Leaf-Spine Architectures:** Leaf-spine networks with many links, especially in dense configurations, are prone to frequent link issues caused by faulty cables or transceivers. In an AI storage fabric, losing even 5% of links between leaf and spine switches can result in a 10–30% performance drop. This is due to the high bandwidth demands, synchronized data flows, and deterministic communication patterns of AI workloads. Link failures disrupt the balanced distribution of traffic, creating congestion on alternate paths, increased latency, and reduced throughput. These disruptions cascade through distributed storage systems, underutilizing compute resources and delaying job completion as stalled GPUs wait for data. The tight coupling between network performance and the AI workload's ability to sustain real-time data delivery amplifies the performance impact.

Optimized Network Support for RDMA

To address the challenges of large-scale AI workloads, AI clusters require advanced network designs that dynamically manage massive data flows, optimize throughput, reduce latency, and maintain network stability. A key standard in these clusters is Remote Direct Memory Access (RDMA), which enables high-performance communication between GPUs by allowing direct data transfers without CPU intervention.

With **NVIDIA GPUDirect**, network adapters and storage drives can directly read from and write to GPU memory, eliminating unnecessary memory copies, reducing CPU overhead, and significantly lowering latency. This enhances performance and efficiency, as RDMA offloads data transfers from the CPU to the GPU or DPU, enabling zero-copy operations that bypass the operating system kernel.

NVIDIA extends these innovations to storage through **GPUDirect Storage**, which establishes a direct data path between GPU memory and local or remote storage—such as NVMe or NVMe over Fabrics (NVMe-oF). By bypassing the CPU's memory buffers, GPUDirect Storage enables a Direct Memory Access (DMA) engine near the NIC or storage device to transfer data directly in and out of GPU memory, further reducing CPU load and enhancing data transfer efficiency. Other high performance AI storage access over NVMe-oF, NFS, S3 over RDMA, or various distributed file systems may also require use of RDMA even if not using GPUDirect Storage.

In AI clusters, RDMA generally requires a lossless network fabric to function effectively, which can be achieved using technologies like RDMA over Converged Ethernet (RoCE). NVIDIA Spectrum-X enhances traditional RoCE with several improvements beyond standard RoCEv2, enabling adaptive routing and sophisticated network management. These innovations ensure efficient data transfer and robust performance, making Spectrum-X ideal for AI workloads, whether the network is dedicated solely to storage or shared with other functions like management traffic.

Typical AI Storage Fabric for Training

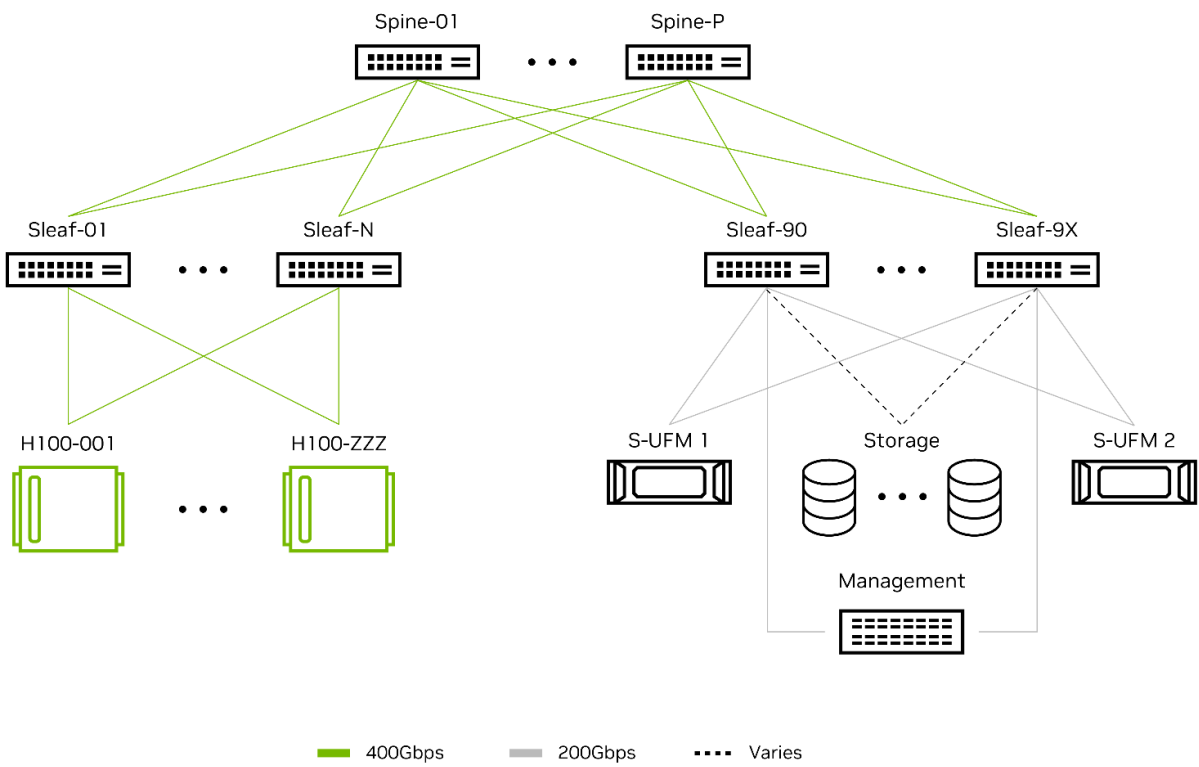
The storage requirements for an AI cluster can vary significantly depending on the specific large language model (LLM) and the size of the dataset. To avoid bottlenecks that could impede GPU performance, certain best practices can guide the setup of an effective AI storage fabric. One key consideration is the ratio of storage nodes to GPU nodes, which ensures sufficient I/O bandwidth.

AI storage vendors generally recommend maintaining a storage bandwidth ratio between GPU nodes and storage nodes in the range of 2:1 to 4:1. This ratio depends on several factors, including GPU generation, network interface cards (NICs), and storage solutions. Following these guidelines helps ensure that storage I/O bandwidth is adequate to keep the GPUs fully utilized and prevents storage from becoming a limiting factor in training performance.

For example, in an NVIDIA DGX H100 SuperPOD configuration, the storage architecture is designed around the concept of Scalable Units (SUs). Each SU typically contains around 32 GPU servers, supported by four storage appliances. In this configuration, each GPU server is equipped with two 400GbE ports for storage connectivity, while each storage appliance has eight 200GbE ports. This configuration creates a storage fabric with 64 x 400GbE ports from the GPU nodes and 32 x 200GbE ports from the storage appliances, delivering the robust and balanced network to handle the I/O demands of large-scale AI workloads. The DGX H100 SuperPOD configuration (shown in the figure below) is based on InfiniBand, but AI Storage Fabrics are also frequently built with Ethernet.

It's important to note: Instead of creating isolated storage fabrics for each SU, best practices recommend a unified storage fabric for the entire AI cluster. This approach is similar to the use of a single east-west (E/W) GPU fabric across the cluster. Since the different GPU nodes in the AI cluster sometimes access storage simultaneously and sometimes access storage at different times, this unified approach ensures efficient resource utilization, simplifies management, and optimizes overall performance for large-scale AI training operations.

Figure 1. Example AI Storage Fabric Based on InfiniBand



NVIDIA Spectrum-X Technology

The NVIDIA Spectrum-X AI networking platform, combining Spectrum Switches and NVIDIA SuperNICs, is the first Ethernet platform designed specifically to improve the performance and efficiency of Ethernet-based AI Clouds. This breakthrough technology achieves higher overall AI performance and energy efficiency, along with consistent and predictable performance in multi-tenant cloud environments.

The major components of the Spectrum-X are described in Table 1.

Table 1. Spectrum-X Components

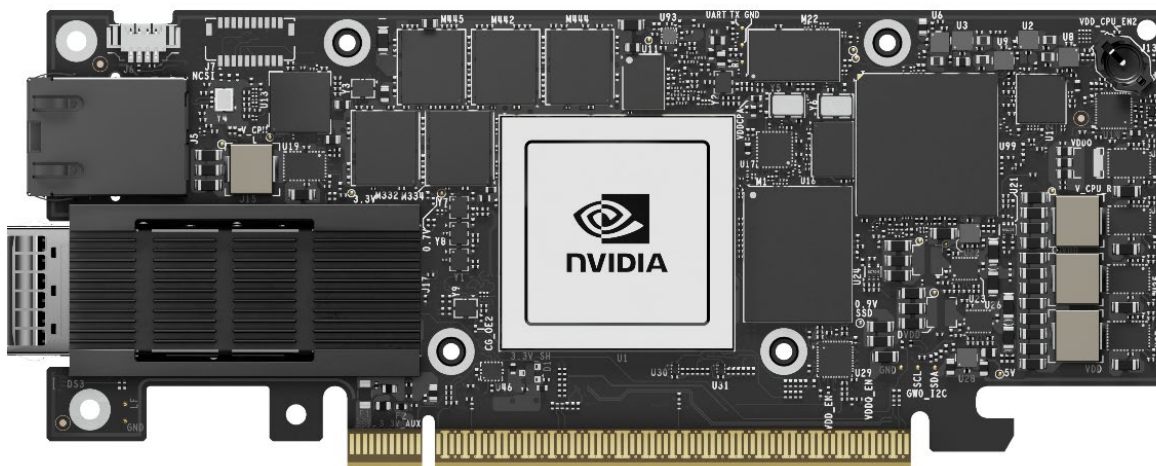
Component	Technology	Description
Compute network	NVIDIA Spectrum-4 SN5600 switch	The SN5600 Ethernet switch provides 64 ports of 800Gbps (OSFP), often can be configured as 128 ports of 400Gbps, offering high port density, low latency, and high performance for AI workloads.
SuperNIC (Compute connectivity)	NVIDIA BlueField-3 or ConnectX-8 SuperNIC	BlueField-3 is purpose-built for AI acceleration, featuring an integrated all-to-all engine for AI, NVIDIA GPUDirect®, and NVIDIA® Magnum IO GPUDirect® Storage (GDS) acceleration technology. It natively runs in network interface card (NIC) mode, leveraging local memory to accelerate large AI clouds. These clouds utilize numerous Queue Pairs, which can be addressed locally within the SuperNIC instead of using server system memory. Lastly, it includes NVIDIA Direct Data Placement (DDP) Technology which augments RoCE adaptive routing. The ConnectX-8 SuperNIC offers similar support for Spectrum-X networking and GDS but with less programmability and fewer storage virtualization offloads than BlueField-3.
Switch OSFP Transceivers	LinkX	800Gbps SR8 switch-side transceiver, supporting two ports of 400GbE Ethernet on multi-mode OM3/OM4 cable.
SuperNIC OSFP Transceivers or splitter cables	LinkX	400Gbps transceiver for the SuperNIC supporting one port of 400Gb/s InfiniBand or 400GbE Ethernet on multi-mode OM3/OM4 cable. 800Gbps Y splitter cable connecting one 800GbE switch port to two 400GbE SuperNIC ports.

400 Gbps optical cables	LinkX®	400GbE/400Gb/s optical cable using 8 data fibers and MPO-12 connectors. Use one for a single switch-to-NIC 400GbE connection or two for an 800GbE switch-to-switch connection.
Network fabric management	NVIDIA NetQ™	The NVIDIA NetQ toolset leverages real-time network telemetry to provide visibility, troubleshooting, and validation of the NVIDIA Spectrum-X fabric, enabling faster location and resolution of problems with a switch, transceiver, or cable.
Network operating system (NOS)	NVIDIA Cumulus Linux®	Cumulus Linux is NVIDIA's flagship network operating system (NOS), bringing innovative features and highest performance and scale to the switches in Spectrum-X networks.
Network simulation	NVIDIA Air	With NVIDIA Air, users can simulate their Spectrum-X fabric in a virtual network, enabling high-fidelity simulation of switches, SuperNICs, and servers. Simulating and testing Spectrum-X configurations, protocols, and changes in a digital environment accelerates Day 0 and Day 1 operations, ensuring faster and smoother production deployment.
SDKs	DOCA™, Magnum IO™, Spectrum SDK/SAI	Spectrum-X supports DOCA, Magnum IO, and Spectrum SDK, providing powerful flexibility and programmability. These tools enable users to tailor their switch and SuperNIC infrastructure to meet their unique requirements and optimize performance.

NVIDIA Spectrum-X SuperNIC

The NVIDIA® Spectrum-X SuperNIC is the 3rd-generation data center infrastructure-on-a-chip that enables organizations to build software-defined, hardware-accelerated IT infrastructures from cloud to core data center to edge. The BlueField®-3 SuperNIC provides up to 400Gb/s of RoCE network connectivity between GPU servers, optimizing peak AI workload efficiency and delivering deterministic and isolated performance between jobs and tenants.

Figure 2. NVIDIA Spectrum-X BlueField-3 Ethernet SuperNIC

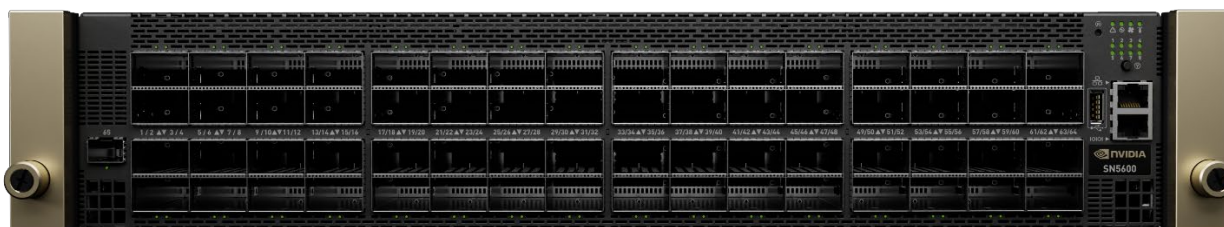


NVIDIA SN5600 Ethernet Switch

Spectrum-X uses the NVIDIA SN5600 Ethernet switches which are built on a 51.2Tbps ASIC, supporting up to 64 ports of 800GbE or 128 ports of 400GbE within a single 2U switch. Setting a new standard, the SN5600 is the first switch designed from the ground up specifically for AI workloads. By integrating specialized high-performance, low latency architecture with standard Ethernet connectivity, the switch delivers unparalleled efficiency and reliability, ensuring seamless operation and optimal performance for demanding AI applications. Spectrum-X switches offer RoCE Extensions for AI with unique enhancements:

- RoCE Adaptive Routing
- RoCE Performance Isolation
- Highest effective bandwidth at scale on standard Ethernet
- Low latency with low jitter and short tail latency

Figure 3. NVIDIA SN5600 Ethernet Switch



Accelerating AI Storage Networks with NVIDIA Spectrum-X

NVIDIA Spectrum-X introduces a game-changing approach to optimizing AI storage networks through its RoCE (RDMA over Converged Ethernet) Adaptive Routing capabilities. Traditional leaf-spine network architectures often struggle with static Equal-Cost Multi-Path (ECMP) routing, which can lead to inefficient load balancing and network congestion. In contrast, Spectrum-X switches use adaptive routing to dynamically manage traffic, ensuring that data flows are distributed across the network in the most efficient manner possible.

RoCE Adaptive Routing enables Spectrum-X switches to select the least congested path within an ECMP group based on real-time congestion metrics. This is achieved by continuously monitoring the egress queue loads on each port, allowing the switch to intelligently forward traffic, on a packet-by-packet basis, to the port with the minimal load. Unlike static routing methods that distribute traffic based on predefined rules, adaptive routing in Spectrum-X is responsive to the current state of the network, balancing loads dynamically to maximize throughput and minimize latency. This ensures that traffic is well-distributed across all available paths, even when dealing with high entropy traffic patterns typical in AI training workloads.

Spectrum-X leverages advanced end-to-end network telemetry gathered from SuperNICs, switches, and optics to detect remote network faults and congestion. This data is utilized by our Global Adaptive Routing technology to make packet forwarding decisions with a fabric-wide perspective, enabling optimal path selection. The result is a higher-performing and more resilient storage fabric.

As packets traverse different network paths, they may arrive at their destination out of order, normally causing the server to discard the packets and request data retransmission. Spectrum-X's integrated SuperNIC technology manages this potential issue seamlessly. At the RoCE transport layer, the SuperNIC directly places packet data directly into memory in the order it was originally sent, regardless of the order in which it was received. This ensures that applications receive data in the correct order. This automatic reordering is invisible to the application, which continues to function as if all packets arrived at the destination sequentially, thereby maintaining data integrity and consistency without additional overhead, retransmission delays, or complexity.

Moreover, Spectrum-X enhances flexibility on the sender side by enabling the SuperNIC to dynamically mark specific traffic for packet reordering eligibility. This ensures that the network enforces inter-packet ordering when necessary, providing both high performance and reliability for AI training applications. Only marked RoCE packets are subject to adaptive routing, providing precise control over which traffic benefits from these advanced routing capabilities.

By leveraging RoCE Adaptive Routing, NVIDIA Spectrum-X optimizes the efficiency and performance of AI storage networks by balancing data flows across the network, reducing congestion, lowering latency, and maximizing network resource utilization. These optimizations are critical for AI clusters requiring rapid, reliable access to vast amounts of data to accelerate training times, and enhance model performance, making Spectrum-X a critical component of the AI infrastructure stack.

RoCE Advanced Congestion Control solves Incast congestion, also known as many-to-one congestion, involving multiple data senders and a single data receiver. This type of congestion cannot be resolved with RoCE Adaptive Routing because all incast traffic eventually flows through a single link. Instead, it requires data-flow metering per endpoint. Spectrum-X RoCE Congestion Control employs end-to-end technology using in-band network telemetry to collect network performance data. The SuperNICs then leverage this collected switch telemetry data to optimally manage and control the data sender's data injection rate, maximizing network sharing efficiency.

Without this congestion control, many-to-one traffic scenarios will cause network back-pressure, congestion spreading, or even packet drops, dramatically degrading network and application performance. In managing this congestion, Spectrum-X SuperNICs execute a congestion control algorithm, capable of handling millions of congestion control events per second with microsecond reaction latency and fine-grained rate adjustments. NVIDIA's congestion control approach significantly improves congestion detection and reaction times by allowing telemetry data to bypass queueing delays in congested flows, while maintaining accurate and concurrent telemetry.

AI Storage Fabric Benchmark Results

Spectrum-X Performance Results from Israel-1

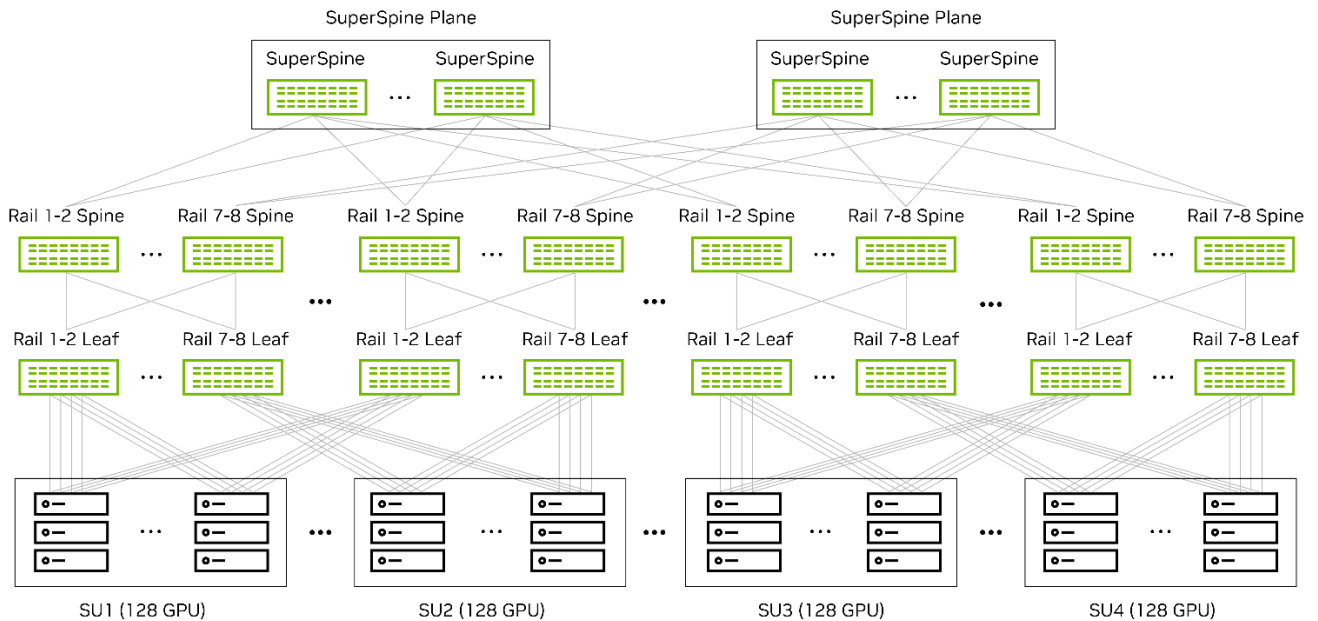
NVIDIA's full stack optimizations deliver significant performance benefits for AI workloads. To fine-tune and optimize Spectrum-X recommended configurations, NVIDIA leverages Israel-1, a generative AI supercomputer and the largest system in Israel.

Figure 4. Israel-1 AI Supercomputer: Validated / Optimized / Quantified

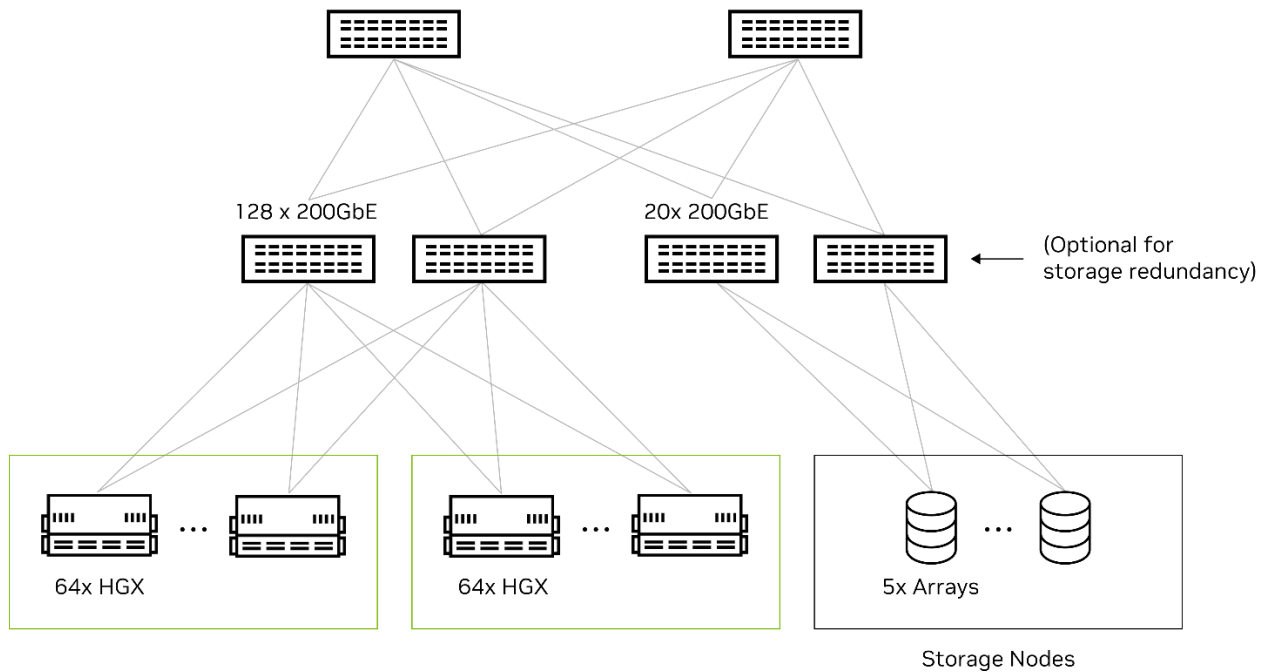


Israel-1 consists of 4 scalable units of NVIDIA Hopper architecture HGX-H100 systems featuring NVIDIA BlueField B3 140 SuperNICs, connected by NVIDIA SN5600 Ethernet switches in a three-tier configuration. Designed as a blueprint for large-scale AI clouds, Israel-1 can be replicated by NVIDIA customers and partners to achieve the same high performance for their own workloads.

Figure 5. Israel-1 SU E/W Topology



In addition to its AI compute capabilities, Israel-1 features 20 all-flash storage servers running the Lustre parallel file system, connected via a dedicated storage fabric. The following benchmarks highlight the AI Storage performance of a fully optimized Spectrum-X network compared to a traditional Ethernet network optimized for AI. While the traditional Ethernet network leverages RoCEv2 and standard Ethernet flow control techniques, it does not include the RoCE Extensions and advanced network optimizations introduced by Spectrum-X.

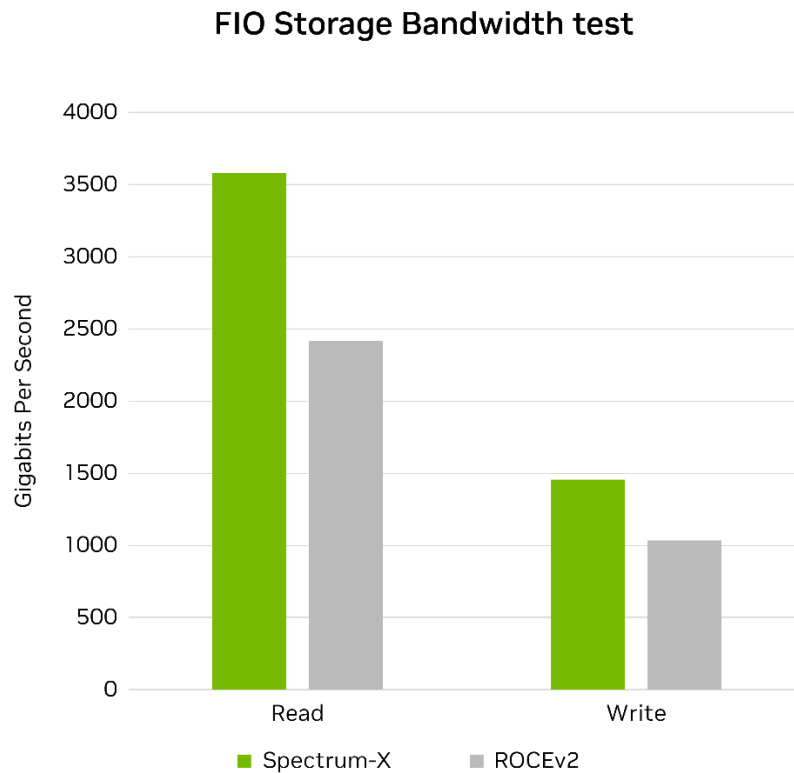
Figure 6. Israel-1 Storage Fabric Topology

To validate the performance benefits of Spectrum-X, we measured the combined bandwidth used by 300 GPUs spread across all 4 Scalable Units in Israel-1. We compared storage performance with the network configured for ROCEv2 against the improved level of performance achieved when the same network was configured to use the Spectrum-X innovations.

The FIO Read bandwidth performance reflects GPU servers requesting data to be read from storage appliances located across the leaf/spine network, similar to many Inference use-cases.

Because Spectrum-X can load balance on a packet-by-packet instead of flow by flow, it avoids ECMP collisions, and deals with congestion in a more proactive way, improving Read performance by up to 48%.

Figure 7. Spectrum-X Increases Storage Bandwidth Through RoCE Adaptive Routing



The FIO Write Bandwidth performance reflects GPU servers sending data to be written to the storage appliances located across the leaf/spine network, like AI Checkpointing. In this case, Spectrum-X improved the Storage write bandwidth by up to 41%.

FIO Storage Benchmark

Flexible I/O Tester (FIO) is a powerful tool for evaluating storage system performance. It simulates various I/O workloads—such as sequential and random read/write operations—to measure metrics like throughput, IOPS, and latency. FIO supports multiple platforms, including Linux, Windows, and macOS) and offers detailed reporting in formats like text, JSON, and CSV. Users can customize test parameters, including block size, queue depth, and access patterns, to match specific needs.

FIO is particularly useful for measuring the performance of storage fabrics attached to AI training clusters. By simulating the high I/O demands typical of AI workloads, FIO helps ensure that the storage infrastructure can meet the performance requirements.

The following command is used to run the FIO benchmark:

```
rm -f /mnt/lustre/sergeya_test/*; fio --filesize=10G -e
xitall_on_error --refill_buffers=1 --zero_buffers=1 --cpus_allowed=56-111 --numa_cpu_nodes=1 --
disk_util=0 --verify=0 --create_serialize=1 --norandommap --loops=1 --
directory=/mnt/lustre/sergeya_test/ --name=file1025.io --ioengine=libaio --randrepeat=0 --exitall -
-direct=1 --numjobs=16 --bs=1m --runtime=60 --ramp_time=2 --iodepth=64 --time_based --
group_reporting --output-format=json --rw=randwrite --nrfiles=1
```

Spectrum-X Improves AI Storage Visibility

NVIDIA NetQ™ is a highly scalable network operations toolset that provides real-time AI network visibility, troubleshooting, and validation. NetQ leverages ‘NVIDIA What Just Happened’ switch telemetry data and NVIDIA® DOCA™ telemetry to deliver actionable insights into the health of Spectrum-X switches and SuperNICs, integrating the network into an organization’s MLOps ecosystem.

NetQ also utilizes flow telemetry to map flow paths and behavior across switch ports and RoCE queues for analyzing specific application flows. NetQ samples packets, analyzes, and reports on per-switch latency (maximum, minimum, and average) and buffer occupancy details along the path of the flow. The NetQ GUI reports all the possible paths, per-path details, and flow behavior. By combining What Just Happened with flow telemetry, NetQ empowers network operators to proactively identify and resolve root causes of server and application issues.

Figure 8. NetQ GUI Reporting



This fabric validation functionality can be integrated into a Storage Vendor's own management tools, offering deep visibility into the state of the AI Storage fabric. NetQ simplifies this process by collecting the relevant data directly from the BlueField-3 adapters, providing a unified connection point for the entire fabric. Alternatively, storage vendors management tools can collect NVIDIA WJH messages directly from the switches using the WJH API for quick and efficient debugging.

Spectrum-X Improves AI Storage Fabric Resiliency

In AI workflows, massive datasets must flow seamlessly between compute and storage nodes to maintain high performance. **Resiliency** in the AI storage network is critical for ensuring reliability, efficiency, and scalability.

1. **Minimized Disruptions:** Network interruptions can delay training, underutilize compute resources, and cause costly re-runs. Resiliency ensures continuous data flow, reducing downtime.

2. **Consistent Performance:** AI training requires predictable low-latency, high-throughput connections. Resilient networks avoid bottlenecks and maintain performance under failure conditions.
3. **Fault Tolerance:** Redundancy and self-healing mechanisms prevent single points of failure, enabling smooth operations even with hardware or link failures.
4. **Data Integrity:** Resilient networks safeguard against packet loss or corruption, ensuring accurate datasets for model training.
5. **Scalability:** As AI workloads grow, resiliency supports larger infrastructures without compromising performance.
6. **Economic Impact:** Downtime or inefficiencies result in wasted resources. Resiliency optimizes uptime, maximizing ROI.
7. **Runs over standard Ethernet:** Non-RDMA traffic runs simultaneously with RDMA traffic, enabling other networked applications to function normally.

The Spectrum-X platform enhances network resiliency through its Global Adaptive Routing feature, which detects remote network faults and remote network congestion events using advanced telemetry function to make adaptive routing decisions.

By continuously monitoring network conditions and adapting routes in real-time, Spectrum-X ensures that AI storage networks remain robust and efficient, even in the face of hardware failures or congestion.

Virtualizing AI Storage Fabrics with Digital Twins in NVIDIA AIR

NVIDIA Air enables the creation of **digital twins** for AI clusters, virtualizing all components of the AI East-West (E/W) network as well as the AI storage fabric, including:

- > **GPU servers:** Simulate compute workloads.
- > **SuperNICs and DPUs:** Enable intelligent networking and offloading.
- > **Switches and cables:** Emulate network topologies.
- > **Storage appliances:** AI data storage and retrieval.

This complete and accurate virtual environment allows for testing, validating, and optimizing the AI infrastructure.

Benefits of Virtualizing the AI Storage System

1. **Risk-Free Testing:** Safely experiment with configurations and workflows.
2. **Faster Deployment:** Pre-validate setups to reduce deployment time.
3. **Cost Efficiency:** Save on physical hardware and lab costs.
4. **Performance Optimization:** Fine-tune the system with simulated workloads.

5. **Training & Collaboration:** Facilitate hands-on learning and streamlined teamwork.
6. **Scalability Planning:** Confidently simulate scale-out scenarios.
7. **Enhanced Troubleshooting:** Reproduce and resolve issues virtually.

Virtualizing AI clusters with NVIDIA Air ensures efficient, secure, and high-performing storage systems designed for AI workloads.

Summary

As AI models become more complex and data-intensive, the limitations of traditional storage networks—such as high latency, network congestion, and inefficient data transfer—can significantly hinder AI training performance. NVIDIA Spectrum-X overcomes these challenges with RoCE Adaptive Routing technology and its integrated SuperNIC, which dynamically balance data flows and minimize bottlenecks, optimizing network throughput and reducing latency.

By ensuring high GPU utilization and streamlining data management, Spectrum-X accelerates training times, enhances AI model development, and improves network resiliency through seamless handling of out-of-order packets, allowing AI workloads to run smoothly and reliably. Integrated into NVIDIA's full-stack AI solution, Spectrum-X offers a cohesive and cost-effective environment for AI development that reduces total cost of ownership while boosting performance.

In conclusion, NVIDIA Spectrum-X transforms AI storage networks by enhancing performance, reducing costs, and improving resiliency, enabling organizations to efficiently develop and deploy AI solutions at scale.

Notice

This document is provided for information purposes only and shall not be regarded as a warranty of a certain functionality, condition, or quality of a product. NVIDIA Corporation ("NVIDIA") makes no representations or warranties, expressed or implied, as to the accuracy or completeness of the information contained in this document and assumes no responsibility for any errors contained herein. NVIDIA shall have no liability for the consequences or use of such information or for any infringement of patents or other rights of third parties that may result from its use. This document is not a commitment to develop, release, or deliver any Material (defined below), 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.

Customer should obtain the latest relevant information before placing orders and should verify that such information is current and complete.

NVIDIA products are sold subject to the NVIDIA standard terms and conditions of sale supplied at the time of order acknowledgement, unless otherwise agreed in an individual sales agreement signed by authorized representatives of NVIDIA and customer ("Terms of Sale"). NVIDIA hereby expressly objects to applying any customer general terms and conditions with regards to the purchase of the NVIDIA product referenced in this document. No contractual obligations are formed either directly or indirectly by this document.

NVIDIA makes no representation or warranty that products based on this document will be suitable for any specified use. Testing of all parameters of each product is not necessarily performed by NVIDIA. It is customer's sole responsibility to evaluate and determine the applicability of any information contained in this document, ensure the product is suitable and fit for the application planned by customer, and perform the necessary testing for the application in order to avoid a default of the application or the product. Weaknesses in customer's product designs may affect the quality and reliability of the NVIDIA product and may result in additional or different conditions and/or requirements beyond those contained in this document. NVIDIA accepts no liability related to any default, damage, costs, or problem which may be based on or attributable to: (i) the use of the NVIDIA product in any manner that is contrary to this document or (ii) customer product designs.

No license, either expressed or implied, is granted under any NVIDIA patent right, copyright, or other NVIDIA intellectual property right under this document. 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 under the patents or other intellectual property rights of the third party, or a license from NVIDIA under the patents or other intellectual property rights of NVIDIA.

Reproduction of information in this document is permissible only if approved in advance by NVIDIA in writing, reproduced without alteration and in full compliance with all applicable export laws and regulations, and accompanied by all associated conditions, limitations, and notices.

THIS DOCUMENT AND ALL NVIDIA DESIGN SPECIFICATIONS, REFERENCE BOARDS, FILES, DRAWINGS, DIAGNOSTICS, LISTS, AND OTHER DOCUMENTS (TOGETHER AND SEPARATELY, "MATERIALS") ARE BEING PROVIDED "AS IS." NVIDIA MAKES NO WARRANTIES, EXPRESSED, IMPLIED, STATUTORY, OR OTHERWISE WITH RESPECT TO THE MATERIALS, AND EXPRESSLY DISCLAIMS ALL IMPLIED WARRANTIES OF NONINFRINGEMENT, MERCHANTABILITY, AND FITNESS FOR A PARTICULAR PURPOSE. TO THE EXTENT NOT PROHIBITED BY LAW, IN NO EVENT WILL NVIDIA BE LIABLE FOR ANY DAMAGES, INCLUDING WITHOUT LIMITATION ANY DIRECT, INDIRECT, SPECIAL, INCIDENTAL, PUNITIVE, OR CONSEQUENTIAL DAMAGES, HOWEVER CAUSED AND REGARDLESS OF THE THEORY OF LIABILITY, ARISING OUT OF ANY USE OF THIS DOCUMENT, EVEN IF NVIDIA HAS BEEN ADVISED OF THE POSSIBILITY OF SUCH DAMAGES. Notwithstanding any damages that customer might incur for any reason whatsoever, NVIDIA's aggregate and cumulative liability towards customer for the products described herein shall be limited in accordance with the Terms of Sale for the product.

Trademarks

NVIDIA, the NVIDIA logo, BlueField, and Spectrum-X 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.

Copyright

© 2025 NVIDIA Corporation. All rights reserved.

