



优化 AI 存储网络架构

NVIDIA Spectrum-X 加速 AI 存储网络

白皮书

目录

- 摘要 1
- AI 集群需要快速存储.....2
 - 慢速 AI 存储的影响.....2
 - AI 存储网络中的挑战.....3
 - 针对 RDMA 优化的网络支持.....4
 - 用于训练的典型 AI 存储网络架构.....4
- NVIDIA Spectrum-X 技术.....6
 - NVIDIA Spectrum-X SuperNIC7
 - NVIDIA SN5600 以太网交换机.....7
- 借助 NVIDIA Spectrum-X 加速 AI 存储网络9
- AI 存储网络架构的基准测试结果..... 11
 - Spectrum-X 性能结果..... 11
 - Spectrum-X 提高 AI 存储可见性 14
 - Spectrum-X 构建高弹性 AI 存储网络..... 15
 - 使用 NVIDIA AIR 中的数字孪生来虚拟化 AI 存储网络架构 16
 - 虚拟化 AI 存储系统的优势..... 16
- 总结..... 18

插图列表

图 1. 基于 InfiniBand 的 AI 存储网络架构示例 5

图 2. NVIDIA Spectrum-X BlueField-3 以太网 SuperNIC..... 7

图 3. NVIDIA SN5600 以太网交换机 8

图 4. AI 集群：经过验证、优化、量化 11

图 5. AI 集群的可扩展单元（SU）东西向（E/W）网络拓扑 12

图 6. AI 集群的存储网络架构拓扑 12

图 7. Spectrum-X 通过 RoCE 动态路由提高存储带宽 13

图 8. NetQ GUI 报告 15

表格列表

表 1. Spectrum-X 组件..... 7

摘要

随着 AI 模型的复杂度和数据需求提高，用户对高性能训练基础设施的需求前所未有。传统存储网络通常难以满足现代 AI 工作负载的海量数据吞吐量和低延迟要求，从而成为训练速度和模型效率的瓶颈。本白皮书探讨了加速存储网络技术在优化 AI 训练集群中的关键作用，并重点介绍 NVIDIA® Spectrum® -X 的功能。

通过实现高速以太网交换、可扩展架构和低延迟数据访问，NVIDIA Spectrum-X 平台提供了一种显著提升存储性能的先进技术解决方案。Spectrum-X 的一个主要特性是其动态路由技术，该技术可将存储读写 I/O 性能提升高达 1.6 倍，从而有效消除 ECMP 冲突。结合其第二个主要特性，即可编程拥塞控制（PCC），这些技术可以协同工作，防止数据瓶颈，并提高整体训练效率。它们还通过加速数据传输和优化检查点来提高 GPU 利用率，同时显著增强 AI 训练集群的弹性。

Spectrum-X 无缝集成到 NVIDIA 的全栈 AI 解决方案中，为 AI 开发和部署提供一致且优化的环境。通过与 NVIDIA 更广泛的 GPU 生态系统、NVIDIA NVLink™ 及 AI 优化库、软件和框架保持一致，Spectrum-X 可确保端到端性能优化并简化基础设施管理。通过详细的案例研究和性能分析，本文展示了 Spectrum-X 对 AI 训练的变革性影响，使其成为希望在 AI 领域保持竞争力的企业不可或缺的关键工具。

AI 集群需要快速存储

在 AI 训练集群中，存储在确保模型训练所需的高效数据流方面发挥着关键作用。AI 训练工作负载通常涉及处理大量数据，必须快速访问和操作这些数据，才能最大化地利用计算资源，尤其是 GPU。更快的存储是至关重要的，它有助于更大限度地减少 GPU 空闲时间，实现持续的数据预提取和缓存，从而为训练流水线提供全新数据。这对于训练大语言模型（LLM）、管理大量数据集，并及时处理批量加载、频繁检查点和日志记录活动是至关重要的。

随着集群规模和训练时间的增长以及故障发生的可能性增加，检查点变得越来越重要。有效的检查点可确保从上次保存的状态恢复训练，从而避免重复计算和时间浪费。此外，快速存储对于分布式训练和并行性是至关重要的，确保多个 GPU 可以稳定的同时处理模型或数据集的不同部分。在这些情况下，慢速存储会产生瓶颈，从而阻碍 AI 工作负载的性能并限制其可扩展性。

AI 推理场景高度依赖于快速存储 IO。每个 GPU 的存储速度会影响检索增强生成式（RAG）AI 工作流中提取阶段的持续时间。存储带宽在键值（Key-Value）缓存中也发挥着关键作用，会影响其效率和可扩展性。

慢速 AI 存储的影响

- **训练瓶颈：**慢速存储通常会迫使 CPU 或 GPU 等待数据读取或写入，从而降低整体训练速度和硬件利用率，这会导致昂贵的计算资源使用效率低下。
- **更长的训练时间：**缓慢的数据加载和低效的检查点保存显著增加了 AI 模型的训练时间。这些延迟会在多次迭代中逐渐累积，导致训练周期延长，从而延迟模型部署。
- **更高的成本：**延长的训练时间直接转化为更高的计算成本，尤其是在按时收费的云环境中，低效的存储会导致不必要的支出，降低 AI 项目的成本效益。
- **实践次数减少：**训练所需的时间增加，使得实践不同的模型或超参数变得更具挑战性。这会减慢研发速度，降低快速创新的能力，并使模型难以适应新数据或不断变化的需求。

通过解决这些存储性能的挑战，AI 训练集群可以实现更高的性能、更低的成本和更大的实践灵活性，最终加速 AI 模型的开发和部署。

AI 存储网络中的挑战

在 AI 集群中，架构通常由通过叶脊（leaf-spine）网络互连的 GPU 节点和存储节点组成。这种网络架构可以是专为存储设计的专用网络架构，也可以是同时管理流量的共享存储网络。网络设置的选择对于优化 AI 训练工作负载的效率和性能至关重要。然而，AI 存储网络面临着一些独特的挑战，这些挑战可能会对其性能和效率产生重大影响，尤其是在依赖传统的叶脊拓扑时。

- 1. 大数据流的负载均衡效率低下：**大象流（GPU 和存储节点之间的大型且持续数据传输）是 AI 训练集群的常见特征。传统的叶脊网络通常使用静态等价多路径（ECMP）路由来均衡多条路径上的各个流量。但是，静态 ECMP 是专为许多小的、短暂的流而设计的，缺乏逐个数据包动态负载均衡流量的能力，并且无法适应 AI 工作负载的不同需求。这会导致负载分配不平衡，这意味着某些路径可能会变得拥塞，而其他路径仍未得到充分利用，从而导致 GPU 和存储节点之间的网络性能降低和数据传输效率低下。
- 2. 吞吐量利用率不理想：**ECMP 的静态特性及其在处理动态和大规模数据流方面的局限性限制了标准叶脊网络的吞吐量。这些网络通常只能实现所有叶脊连接理论总吞吐量的 60% 左右，远远低于 AI 工作负载所需的性能。高吞吐量对于更大限度地提高 GPU 利用率和缩短训练时间至关重要。在叶脊网络还能处理带内管理网络上的管理流量的情况下，带宽竞争会进一步降低可用于存储操作的有效吞吐量。
- 3. 交换机间的拥塞和高延迟：**由于静态负载均衡导致网络路径利用不均匀，会导致交换机间的拥塞，这不仅降低了吞吐量，还增加了延迟。在 AI 集群中，高延迟会破坏高效训练所需的数据流畅通，迫使 GPU 闲置等待数据，并降低整体训练速度。在结合了存储和管理流量的网络中，这一问题将进一步加剧，高延迟的情况也将变得更加频繁。
- 4. 密集叶脊架构中频繁出现的链路问题：**具有许多链路的叶脊网络（尤其是在密集配置中）容易因线缆或收发器故障而频繁出现链路问题。在 AI 存储网络架构中，叶脊交换机之间丢失 5% 的链路，会导致性能下降 10 - 30%。这是由于 AI 工作负载的高带宽需求、同步数据流和确定性通信模式。链路故障会破坏流量的均衡分布，造成备用路径的拥塞，增加延迟并降低吞吐量。这些中断会通过分布式存储系统传播，并由于停滞的 GPU 在等待数据，导致计算资源利用不足和作业延迟完成。网络性能与 AI 工作负载维持实时数据传输的能力之间的紧密耦合放大了对性能影响。

针对 RDMA 优化的网络支持

为了应对大规模 AI 工作负载的挑战，AI 集群需要先进的网络设计来动态管理海量数据流、优化吞吐量、降低延迟并保持网络稳定性。这些集群中的一个关键标准是远程直接内存访问（RDMA），它允许进行直接数据传输而无需 CPU 干预，从而实现 GPU 之间的高性能通信。

借助 NVIDIA GPUDirect，网卡和存储驱动器可以直接读取和写入 GPU 内存，从而消除不必要的内存复制、减少 CPU 开销并显著降低延迟。这提高了性能和效率，因为 RDMA 将数据传输从 CPU 卸载到 GPU 或 DPU，从而实现绕过操作系统内核的零复制操作。

NVIDIA 通过 GPUDirect Storage 将这些创新扩展到存储领域，GPUDirect Storage 在 GPU 内存与本地或远程存储（例如 NVMe 或 NVMe over Fabric（NVMe-oF））之间建立直接数据路径。通过绕过 CPU 的内存缓冲区，GPUDirect Storage 使靠近网卡或存储设备的直接内存访问（DMA）引擎能够直接将数据传入和传出 GPU 内存，从而进一步降低 CPU 负载并提高数据传输效率。通过 NVMe-oF、NFS、S3 over RDMA 或各种分布式文件系统也可进行的其他高性能 AI 存储访问，即使不使用 GPUDirect Storage，也能够使用 RDMA。

在 AI 集群中，RDMA 通常需要无损网络架构才能有效运行，这可以使用 RDMA over Converged Ethernet（RoCE）等技术来实现。NVIDIA Spectrum-X 在标准 RoCEv2 的基础上进行了多项改进，增强了传统 RoCE，实现了动态路由和复杂的网络管理。这些创新可确保高效的数据传输和强大的性能，使 Spectrum-X 成为 AI 工作负载的理想选择，无论网络是专用于存储，还是与管理流量等其他功能共享。

用于训练的典型 AI 存储网络架构

AI 集群的存储需求因特定的大语言模型（LLM）和数据集大小有可能存在显著差异。为避免出现 GPU 性能瓶颈，设置合适的存储网络架构至关重要。一个关键考虑因素是存储节点与 GPU 节点的带宽比例，以确保数据传输速度足够支持 GPU 处理高并发 I/O 任务。

AI 存储供应商通常建议将 GPU 节点与存储节点之间的存储带宽比保持在 2:1 到 4:1 的范围内。此比率取决于多个因素，包括 GPU 型号、网卡（NIC）性能和存储解决方案的类型。遵循这些准则有助于确保存储 I/O 带宽足以充分利用 GPU，并防止存储成为训练性能的限制因素。

例如，在 NVIDIA DGX Hopper GPU SuperPOD 配置中，存储架构是围绕可扩展单元（SU）的概念设计的。每个 SU 通常包含大约 32 台 GPU 服务器，并由四台存储设备提供支持。在此配置中，每台 GPU 服务器配有两个 400GbE 端口用于存储连接，而每台存储设备都有 8 个 200GbE 端口。此配置创建了一个存储网络架构，其中包含来自 GPU 节点的 64 个 400GbE 端

口和来自存储设备的 32 个 200GbE 端口，从而提供强大而均衡的网络来处理大规模 AI 工作负载的 I/O 需求。DGX Hopper GPU SuperPOD 配置（如下图所示）是基于 InfiniBand 的，但 AI 存储网络架构也经常使用以太网来构建。

值得注意的是：最佳实践建议为整个 AI 集群使用统一的存储网络架构，而不是为每个 SU 创建独立的存储网络架构。这种方法类似于在整个集群中使用单个东西向（E/W）GPU 网络架构。由于 AI 集群中的不同 GPU 节点有时会同时访问存储，有时会在不同时间访问存储，因此这种统一的方法可确保高效的资源利用，简化管理，并优化大规模 AI 训练操作的整体性能。

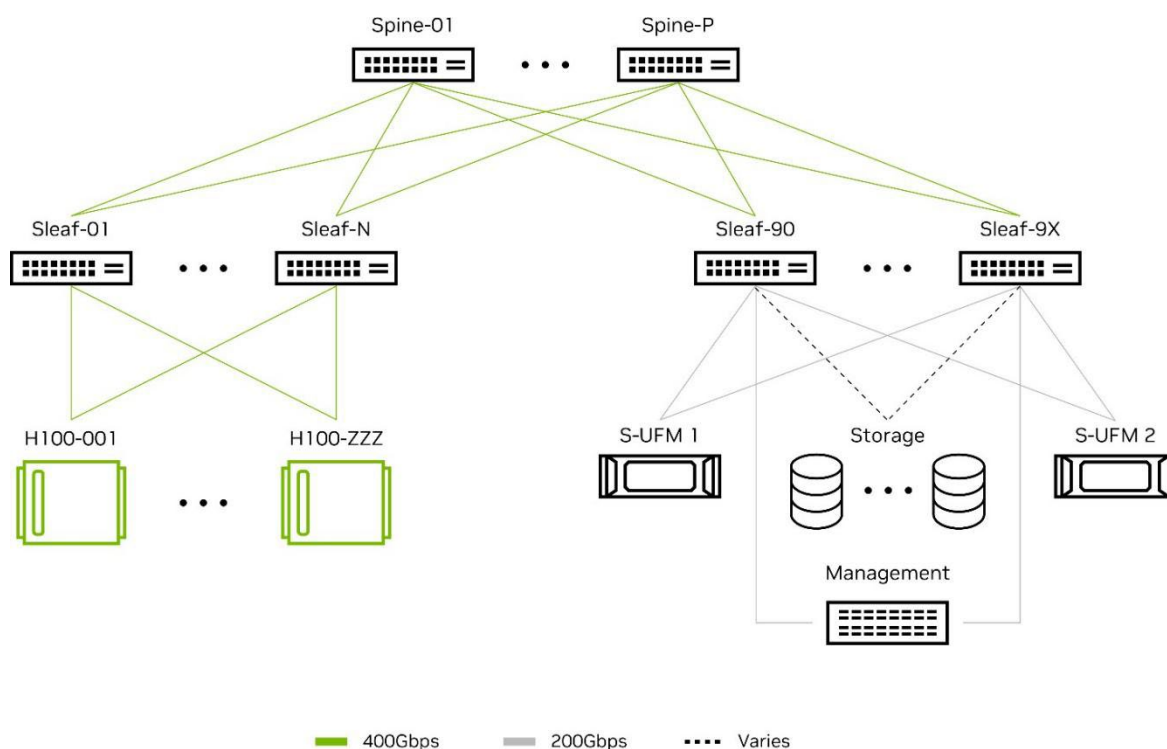


图 1. 基于 InfiniBand 的 AI 存储网络架构示例

NVIDIA Spectrum-X 技术

NVIDIA Spectrum-X AI 网络平台结合了 Spectrum 交换机和 NVIDIA SuperNIC，是首个专为提高基于以太网的 AI 云性能和效率而设计的以太网平台。这项突破性技术可实现更高的整体 AI 性能和能效，并在多租户云环境中实现一致且可预测的性能。

Spectrum-X 的主要组件如表 1 所示。

组件	技术	说明
计算网络	NVIDIA Spectrum-4 SN5600 交换机	SN5600 以太网交换机提供 64 个 800Gbps（OSFP）端口，通常可配置为 128 个 400Gbps 端口，为 AI 工作负载提供高端口密度、低延迟和高性能。
SuperNIC（计算连接）	NVIDIA BlueField-3 或 ConnectX-8 SuperNIC	BlueField-3 专为 AI 加速而设计，具有用于 AI 的集成 all-to-all 引擎、NVIDIA GPUDirect® 和 NVIDIA® Magnum IO GPUDirect® Storage（GDS）加速技术。它以网卡（NIC）模式在本地运行，利用本地内存加速大型 AI 云。这些云使用大量队列对，这些队列对可以在 SuperNIC 内进行本地寻址，而无需使用服务器系统内存。最后，它包括 NVIDIA 直接数据放置（DDP）技术，可增强 RoCE 动态路由功能。ConnectX-8 SuperNIC 为 Spectrum-X 网络和 GDS 提供了类似的支持，但与 BlueField-3 相比，它的可编程性低，存储虚拟化卸载有限。
交换机 OSFP 收发器	LinkX	800Gbps SR8 交换机侧收发器，在多模 OM3/OM4 线缆上支持两个 400GbE 以太网端口。
SuperNIC OSFP 收发器或分线线缆	LinkX	适用于 SuperNIC 的 400Gbps 收发器，通过多模 OM3/OM4 线缆支持单个 400Gb/s InfiniBand 或 400GbE 以太网端口。800Gbps Y 型分线线可将一个 800GbE 交换机端口连接到两个 400GbE SuperNIC 端口。
400 Gbps 光缆	LinkX®	400GbE/400Gb/s 光缆使用 8 条数据光纤和 MPO-12 连接器。使用一根光缆实现 400GbE 交换机到网卡的连接，或使用两根光缆实现 800GbE 交换机到交换机的连接。
网络架构管理	NVIDIA NetQ™	NVIDIA NetQ 工具集利用实时网络遥测技术提供 NVIDIA Spectrum-X 网络架构的可视化、故障排除和验证，从而能够更快地定位和解决交换机、收发器或线缆的问题。
网络操作系统（NOS）	NVIDIA Cumulus Linux®	Cumulus Linux 是 NVIDIA 的旗舰网络操作系统（NOS），可为 Spectrum-X 网络中的交换机带来创新功能、超高性能和可扩展性。

网络仿真	NVIDIA Air	借助 NVIDIA Air，用户可以在虚拟网络中仿真其 Spectrum-X 网络架构，从而实现交换机、SuperNIC 和服务器的仿真。在数字环境中仿真和测试 Spectrum-X 配置、协议和更改，可加速第 0 天和第 1 天的运营，确保更快、更流畅的生产部署。
SDK	DOCA™、Magnum IO™、Spectrum SDK/SAI	Spectrum-X 支持 DOCA、Magnum IO 和 Spectrum SDK，具有强大的灵活性和可编程性。这些工具使用户能够定制其交换机和 SuperNIC 基础设施，以满足其独特需求并优化性能。

表 1. Spectrum-X 组件

NVIDIA Spectrum-X SuperNIC

NVIDIA® Spectrum-X SuperNIC 是第三代片上数据中心基础设施，可帮助企业构建从云到核心数据中心再到边缘的软件定义、硬件加速的 IT 基础设施。BlueField® -3 SuperNIC 可在 GPU 服务器之间提供高达 400Gb/s 的 RoCE 网络连接，优化峰值 AI 工作负载效率，并在作业和租户之间提供确定且隔离的性能。

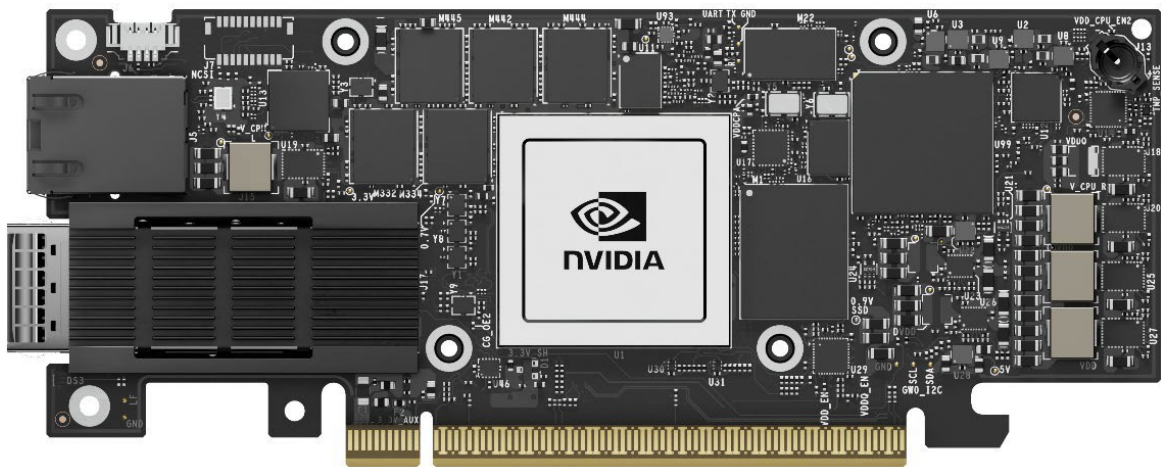


图 2. NVIDIA Spectrum-X BlueField-3 以太网 SuperNIC

NVIDIA SN5600 以太网交换机

Spectrum-X 使用基于 51.2 Tbps ASIC 构建的 NVIDIA SN5600 以太网交换机，在单个 2U 交换机内支持多达 64 个 800GbE 端口或 128 个 400GbE 端口。

SN5600 是首款完全专为 AI 工作负载而设计的交换机，树立了全新的标准。通过将高性能、低延迟的专用架构与标准以太网连接相集成，该交换机可提供无与伦比的效率和可靠性，确保为要求苛刻的 AI 应用程序提供无缝运营和卓越性能。Spectrum-X 交换机为 AI 提供具有独特增

强功能的 RoCE 扩展：

- > RoCE 动态路由
- > RoCE 性能隔离
- > 在标准以太网上实现大规模的更高有效带宽
- > 具有低抖动和短尾部延迟的低时延



图 3. NVIDIA SN5600 以太网交换机

借助 NVIDIA Spectrum-X 加速 AI 存储网络

NVIDIA Spectrum-X 引入了一种颠覆性的方法，即通过其 RoCE（RDMA over Converged Ethernet）动态路由功能来优化 AI 存储网络。传统的叶脊网络架构通常难以应对静态的等价多路径（ECMP）路由，这可能导致负载均衡的效率低下和网络拥塞。相比之下，Spectrum-X 交换机使用动态路由来动态管理流量，确保以尽可能高效的方式在网络中分发数据流。

RoCE 动态路由使 Spectrum-X 交换机能够根据实时拥塞指标选择 ECMP 组中拥塞最轻的路径。这是通过持续监控每个端口上的出口队列负载来实现的，允许交换机以逐包的方式智能转发流量到负载最少的端口。与基于预定义规则分配流量的静态路由方法不同，Spectrum-X 中的动态路由能够响应网络的当前状态，动态均衡负载以更大限度地提高吞吐量并更大限度地减少延迟。这可确保流量在所有可用路径中得到均匀分发，即使在处理 AI 训练工作负载中常见的高熵流量模式时也是如此。

Spectrum-X 利用从 SuperNIC、交换机和光学器件收集的先进端到端网络遥测数据来检测远程网络故障和拥塞。我们的全局动态路由技术利用这些数据从整个网络架构的角度做出数据包转发决策，从而实现最佳路径选择。从而获得性能更高、弹性更强的存储网络架构。

当数据包经过不同的网络路径时，它们可能会无序到达目的地，这通常会导致服务器丢弃数据包并请求重新传输数据。Spectrum-X 集成的 SuperNIC 技术可以无缝管理这一潜在问题。在 RoCE 传输层，SuperNIC 直接将数据按照最初发送的顺序放置于内存，而不管其接收的顺序如何。这可确保应用程序以正确的顺序接收数据。这种自动重新排序对于应用程序是不可见的，应用程序将继续运行，就像所有数据包都是按顺序到达目的地一样，从而保持数据的完整性和一致性，而不会产生额外的开销、重传延迟或复杂性。

此外，Spectrum-X 通过支持 SuperNIC 动态标记特定流量以符合数据包重新排序条件，增强了发送方的灵活性。这可确保网络在必要时强制执行数据包间排序，从而为 AI 训练应用程序提供高性能和可靠性。只有标记的 RoCE 数据包才能进行动态路由，从而精确控制哪些流量可以从这些高级路由功能中受益。

通过利用 RoCE 动态路由，NVIDIA Spectrum-X 可以均衡整个网络的数据流、减少拥塞、降低延迟并更大限度地提高网络资源利用率，从而优化 AI 存储网络的效率和性能。这些优化对于需要快速、可靠地访问大量数据以加快训练速度并增强模型性能的 AI 集群至关重要，这使得 Spectrum-X 成为 AI 基础设施堆栈的关键组件。

RoCE 高级拥塞控制解决了涉及多个数据发送方和单个数据接收方的 Incast 拥塞（也称为多打一拥塞）问题。这种类型的拥塞无法通过 RoCE 动态路由解决，因为所有 Incast 流量最终都会流经单个链路。相反，它需要对每个端点进行数据流测量。Spectrum-X RoCE 拥塞控制采用端到端技术，使用带内网络遥测来收集网络性能数据。然后，SuperNIC 利用收集到的交换机遥测数据来优化管理和控制数据发送方的数据注入速率，从而更大限度地提高网络共享效率。

如果没有这种拥塞控制，多对一流量场景将导致网络背压、拥塞扩散甚至数据包丢失，从而严重降低网络和应用程序的性能。在管理这种拥塞时，Spectrum-X SuperNIC 会执行拥塞控制算法，该算法能够在微秒级的反应延迟和细粒度的速率调整下，实现每秒处理数百万个拥塞控制事件。NVIDIA 的拥塞控制方法允许遥测数据绕过拥塞流中的队列延迟，同时保持准确和并发的遥测，从而显著改善拥塞检测和反应时间。

AI 存储网络架构的基准测试结果

Spectrum-X 性能结果

NVIDIA 的全栈优化为 AI 工作负载带来了显著的性能优势。为了微调和优化 Spectrum-X 的推荐配置，NVIDIA 利用了一个大型生成式 AI 集群。



图 4. AI 集群：经过验证、优化、量化

此 AI 集群由 4 个 NVIDIA Hopper 架构 HGX 系统的可扩展单元（SU）组成，这些 HGX 系统配备了 NVIDIA BlueField B3140 SuperNIC，并通过三层配置的 NVIDIA SN5600 以太网交换机进行连接。此 AI 集群是为大规模 AI 云为目标而设计的，可供 NVIDIA 客户和合作伙伴来复制该设计，从而使他们自己的工作负载实现相同的高性能。

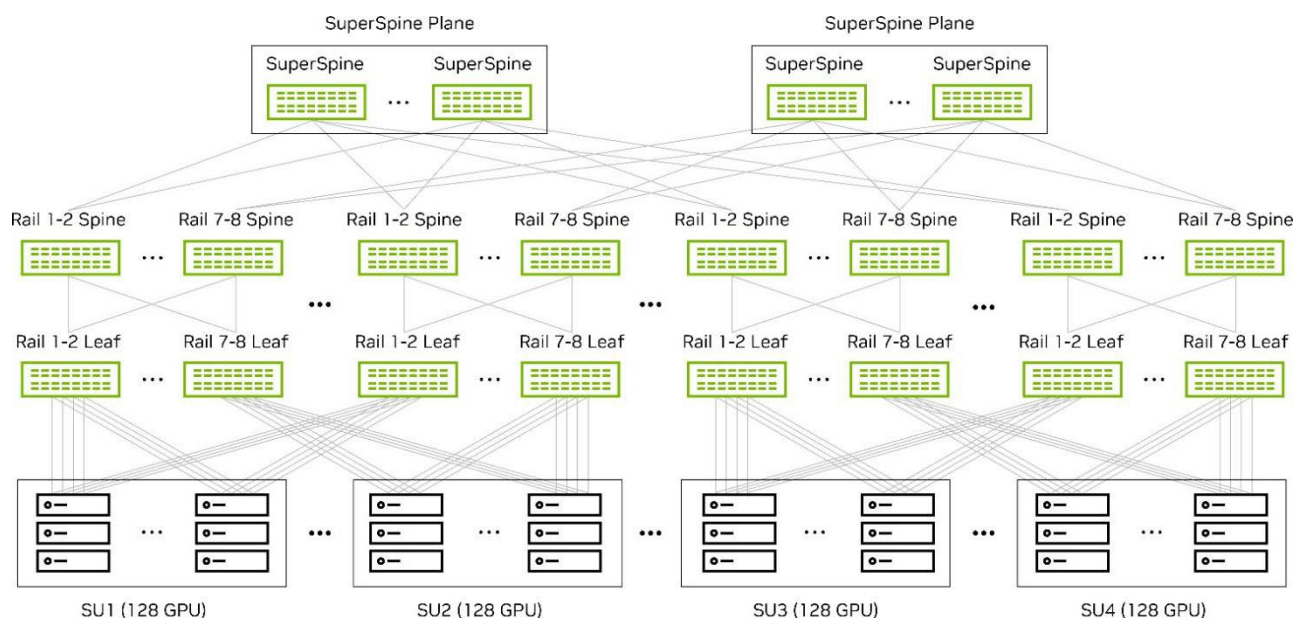


图 5. AI 集群的可扩展单元 (SU) 东西向 (E/W) 网络拓扑

除了 AI 计算能力之外，此 AI 集群还配备 20 台运行 Lustre 并行文件系统的全闪存存储服务器，并通过专用存储网络架构进行连接。以下基准测试重点介绍了完全优化的 Spectrum-X 网络与针对 AI 优化的传统以太网网络相比的 AI 存储性能。虽然传统以太网网络利用 RoCEv2 和标准以太网流量控制技术，但它不包括 Spectrum-X 引入的 RoCE 扩展和高级网络优化。

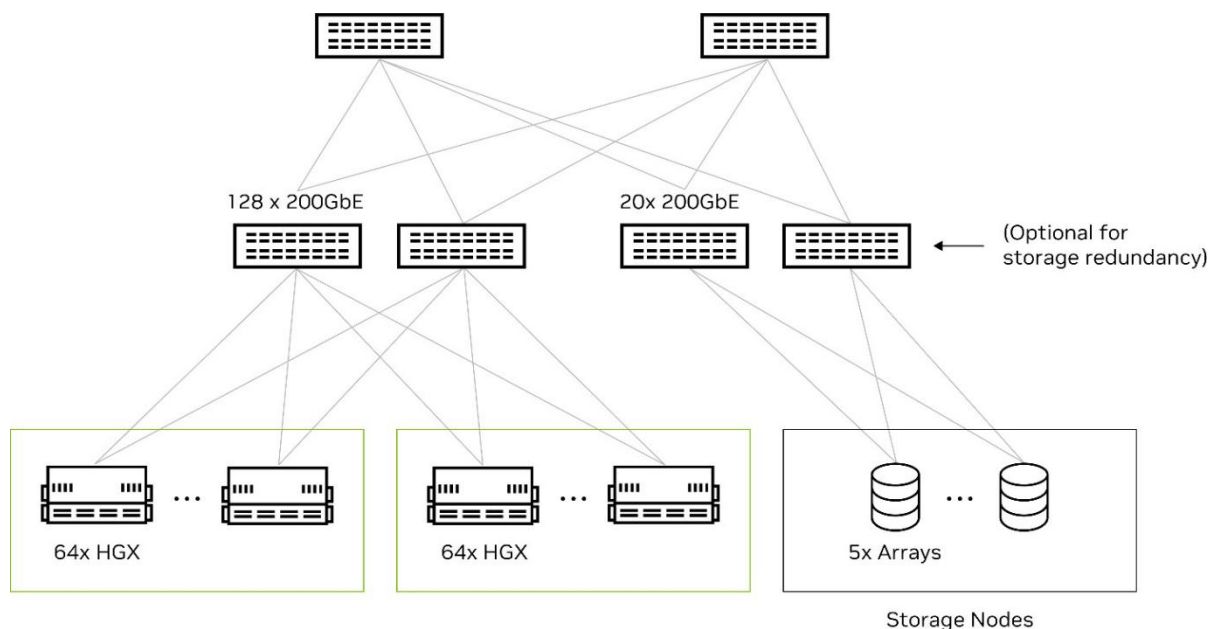


图 6. AI 集群的存储网络架构拓扑

为了验证 Spectrum-X 的性能优势，我们测量了此 AI 集群所有 4 个可扩展单元中 300 个 GPU 的总带宽。我们将适用于 RoCEv2 的网络配置与使用 Spectrum-X 创新实现更高性能水平的同一网络配置进行了比较。FIO 读取带宽性能反映了 GPU 服务器请求从位于叶脊网络中的存储设备读取数据的情况，类似于许多推理用例。由于 Spectrum-X 可以逐包而不是逐流进行负载均衡，因此它可以避免 ECMP 冲突，并以更主动的方式处理拥塞，从而将读取性能提高多达 48%。

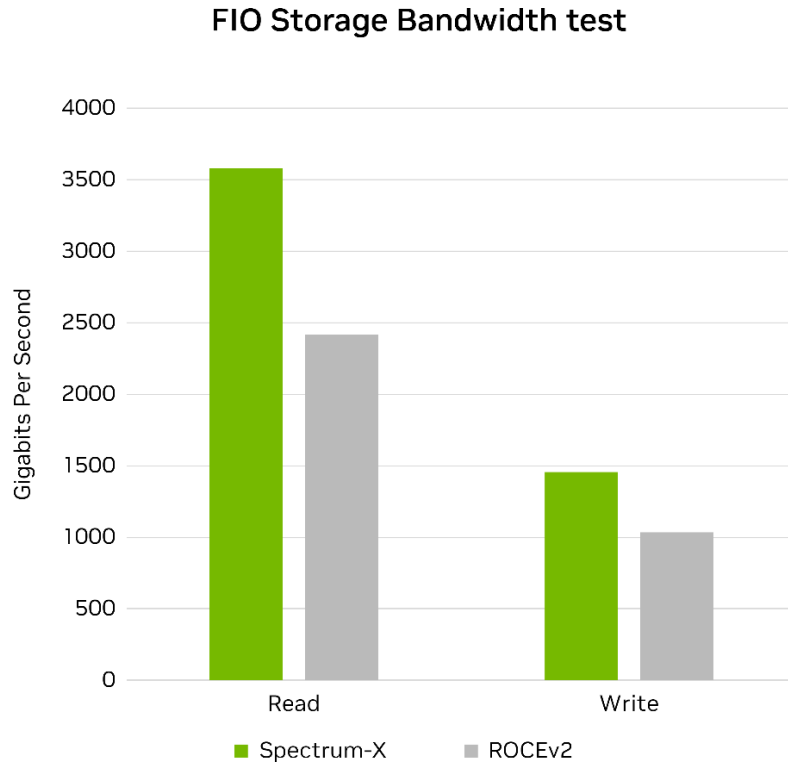


图 7. Spectrum-X 通过 RoCE 动态路由提高存储带宽

FIO 写入带宽性能反映了 GPU 服务器将要写入位于叶脊网络中存储设备的数据（例如 AI Checkpointing）发送至存储设备的情况。

在这种情况下，Spectrum-X 将存储写入带宽提高了 41%。

FIO 存储基准测试

灵活的 I/O 测试器（FIO）是评估存储系统性能的强大工具。它模拟各种 I/O 工作负载（例如顺序和随机读/写操作），以测量吞吐量、IOPS 和延迟等指标。FIO 支持多个平台（包括 Linux、Windows 和 macOS），并以文本、JSON 和 CSV 等格式提供详细报告。用户可以自定义测试参数（包括块大小、队列深度和访问模式），以满足特定需求。

FIO 对于测量连接到 AI 训练集群的存储网络架构的性能特别有用。通过模拟 AI 工作负载典型的

高 I/O 需求，FIO 有助于确保存储基础设施能够满足性能需求。

以下命令用于运行 FIO 基准测试：

```
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 提高 AI 存储可见性

NVIDIA NetQ™ 是一款高度可扩展的网络运营工具集，可提供实时 AI 网络可视化、故障排除和验证。NetQ 利用 “NVIDIA What Just Happened (WJH)” 交换机遥测数据和 NVIDIA® DOCA™ 遥测技术，提供有关 Spectrum-X 交换机和 SuperNIC 运行状况的可行见解，将网络集成到企业的 MLOps 生态系统中。

NetQ 还利用流遥测技术映射交换机端口和 RoCE 队列之间的流路径和行为，以分析特定的应用程序流。NetQ 会对数据包进行采样、分析并报告每台交换机的延迟（最大、最小和平均），以及流路径中缓冲区占用情况的详细信息。NetQ GUI 可报告所有可能的路径、每个路径的详细信息和流行为。通过将 What Just Happened 与流遥测相结合，NetQ 使网络运营商能够主动识别并解决服务器和应用程序的问题。



图 8. NetQ GUI 报告

这种网络架构验证功能可以集成到存储供应商自己的管理工具中，从而提供对 AI 存储网络架构状态的深度可视化。NetQ 通过直接从 BlueField-3 网卡收集相关数据来简化此过程，为整个网络架构提供统一的连接点。除此之外，存储供应商的管理工具可以使用 WJH API 直接从交换机收集 NVIDIA WJH 消息，以便快速高效地进行问题的诊断。

Spectrum-X 构建高弹性 AI 存储网络

在 AI 工作流程中，海量数据集需要在计算和存储节点之间无缝、高效地流动，这是维持高性能的关键。因此，AI 存储网络的弹性对于保障可靠性、效率和可扩展性至关重要。

1. **最小化中断：**网络中断会导致训练延迟、计算资源利用率低下，甚至可能需要付出高昂代价重新训练。高弹性网络可确保数据持续流动，从而最大限度减少停机时间。

2. **性能一致性：**AI 训练需要可预测的低延迟、高吞吐量连接。弹性网络可避免出现瓶颈，并在故障条件下保持性能稳定。
3. **高容错性：**通过冗余和自愈机制，弹性网络能够防止单点故障，即使在硬件或链路发生故障时也能确保系统平稳运行。
4. **数据完整性：**弹性网络可防止数据包丢失或损坏，确保模型训练的数据集准确。
5. **可扩展性：**随着 AI 工作负载的增长，弹性网络可支持更大的基础设施，且不会影响性能。
6. **经济效益：**停机或低效运行会导致资源浪费。弹性网络可优化实现正常运行时间，从而更大限度地提高投资回报率（ROI）。
7. **标准以太网兼容性：**支持非 RDMA 流量与 RDMA 流量同时运行，确保其他网络应用程序正常运行。

Spectrum-X 平台通过其全局动态路由功能显著增强了网络弹性，该功能利用高级遥测技术检测远程网络故障和拥塞事件，从而做出智能的动态路由决策。

通过持续监控网络状况并实时调整路由策略，Spectrum-X 可确保 AI 存储网络即使出现硬件故障或拥塞的情况下也能保持稳健和高效。

使用 NVIDIA AIR 中的数字孪生来虚拟化 AI 存储网络架构

NVIDIA Air 提供了为 AI 集群创建数字孪生的能力，可以虚拟化 AI 东西向（E/W）网络以及 AI 存储网络架构的所有组件，包括：

- > **GPU 服务器：**模拟计算工作负载。
- > **SuperNIC 和 DPU：**实现智能网络和卸载。
- > **交换机和线缆：**模拟网络拓扑。
- > **存储设备：**AI 数据存储和检索。

这一完整且精确的虚拟环境可用于 AI 基础设施的测试、验证和优化。

虚拟化 AI 存储系统的优势

1. **无风险测试：**在安全的环境中试验各种配置和工作流程。
2. **更快的部署：**通过预先验证配置来缩短部署时间。

3. **成本效益：**节省物理硬件和实验成本。
4. **性能优化：**使用模拟工作负载对系统进行精细化调整。
5. **培训和协作：**促进实践学习和简化团队合作。
6. **可扩展性规划：**自信地模拟扩展场景。
7. **增强的故障排除：**通过虚拟环境重现并解决问题。

使用 NVIDIA Air 虚拟化 AI 集群可保证为 AI 工作负载高效的设计安全且高性能的存储系统。

总结

随着 AI 模型日益复杂和数据密集化，传统存储网络高延迟、网络拥塞和数据传输效率低下等方面的局限性，严重制约 AI 训练性能。NVIDIA Spectrum-X 凭借 RoCE 动态路由技术及其集成的 SuperNIC 克服了这些挑战，这些技术可动态均衡数据流并更大限度地减少瓶颈，从而优化网络吞吐量、降低延迟。

通过确保 GPU 的高利用率和简化数据管理，Spectrum-X 可通过无缝处理乱序数据包来加快训练时间、增强 AI 模型开发并提高网络弹性，从而使 AI 工作负载平稳可靠地运行。作为 NVIDIA 全栈 AI 解决方案的一部分，Spectrum-X 为 AI 开发提供了一个高性能、高性价比的环境，在提升性能的同时降低了总体拥有成本（TCO）。

总而言之，NVIDIA Spectrum-X 通过增强性能、降低成本和提高弹性来改变 AI 存储网络，助力企业高效地进行大规模 AI 解决方案的开发与部署。

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.

