

超节点与 Scale up 网络专题之英伟达：行业标杆，领先优势建立在 NVLink 和 NVLink Switch

2026 年 2 月 5 日

看好/维持

通信

行业报告

分析师

石伟晶 电话：021-25102907 邮箱：shi_wj@dxzq.net.cn

执业证书编号：S1480518080001

投资摘要：

大语言模型（LLM）参数规模从千亿级向万亿级乃至十万亿级演进，跨服务器张量并行（TP）成为必然选择；此外混合专家（MoE）模型在 Transformer 架构 LLM 中的规模化应用，更使跨服务器专家并行（EP）成为分布式训练和推理的关键技术需求。为应对 TP 和 EP 对网络带宽与延迟的极为严苛的要求，构建超带宽、超低延迟的 Scale up 网络（纵向扩张网络）成为业界主流技术路径。

目前英伟达超节点已经推出成熟方案。2024-2026 年，英伟达陆续推出 GH200 NVL72、GB200/ GB300 NVL72、VR200 NVL72 三代超节点。

- **Hopper 架构开启超节点 Scale up 初步探索。**GH200 通过 NVLink 和 NVLink-C2C（Chip-to-Chip）技术，使得每个 GPU 可以访问其他所有 CPU 和 GPU 芯片的内存，实现 GPU 与 CPU 内存统一编址。
- **Blackwell 架构推动 Scale up 标准化。**GB200 NVL72 将 Scale-up 规模稳定在 72 个 GPU/机柜，形成可复制标准化方案。NVL72 由 18 个 Compute Tray（计算托架）和 9 个 Switch Tray（网络交换托架）构成。其中，Compute Tray 是计算核心单元，负责提供强大的计算能力；Switch Tray 是高速通信枢纽，用于实现 GPU 之间的高速数据交换。NVL72 背板通过“NVLink5 私有协议 + 铜线缆”将 18 个 Compute Tray 中的 72 颗 B200 GPU 和 9 个 Switch Tray 中的 18 颗 NVSwitch 芯片进行满带宽全连接。
- **Rubin 架构推动 Scale up 方案带宽倍增。**2026 年 1 月 CES 展会，英伟达发布 Rubin 架构 VR200 NVL72。其中 NVLink 6 Switch 实现单 GPU 的互连带宽提升至 3.6 TB/s，上代为 1.8TB/s。Scale out 方面，Spectrum-6 交换机支持 CPO（共封装光学）技术，将 32 个 1.6Tb/s 硅光光学引擎与交换芯片直接封装集成。

在超节点方案上，英伟达处于领先优势。2024-2025 年，英伟达陆续推出 GH200 NVL72、GB200/ GB300 NVL72 等成熟超节点解决方案。根据大摩预测，2025 年英伟达 GB200/300 NVL72 出货量约 2800 台。展望 2026-2027 年，英伟达计划推出 Vera Rubin NVL144 和 Rubin Ultra NVL576。互联 GPU 数将从 72 颗进一步向 576 颗发展。届时，英伟达将在新一代 Kyber 机架架构中引入 NVLink Switch Blade（NVLink 交换机刀片），通过 PCB 中板替代传统 5000+根有源铜缆。可以看到，Rubin Ultra NVL576 仍具有较强的工程创新能力。

英伟达超节点的优势建立在 NVLink 和 NVLink Switch。为实现 AI 训练集群高带宽与低延迟数据传输，NVLink 重新设计通信架构，并引入一系列先进技术，包括网状拓扑、差分信号传输、流量调度信用机制、多 Lane 绑定技术、统一内存空间等。截止 2025 年，NVLink 5 Switch 实现支持单 GPU 到 GPU 带宽 1800GB/s，可构建 72 GPU 的 NVLink 域，总带宽达 130 TB/s（双向），支持 72 GPU 全互联通信。在后续计划中，NVSwitch Gen6 和 Gen7 的 GPU-to-GPU 通信带宽继续升级为 3.6TB/s。

但另一方面，Scale up 网络兴起源于满足大模型分布式训练和推理中的张量并行(TP)与专家并行(EP)。目前 AI 产业也在探索降低 TP 与 EP 规模的技术方案，从而降低 Scale up 网络规模的上限。我们认为，Scale up 网络的发展空间或限制英伟达在超节点领域的领先优势。为保持领先优势，实现 Scale up 网络和 Scale out 网络融合或将成为英伟达超节点新的发展趋势。

投资策略：

自 2025 年开始，超节点成为 AI 算力网络重要的技术创新方向。从 AI 基建竞争维度，AI 芯片厂商从芯片算力性能竞争延续至芯片+Scale up 网络的双战场。因此，除了原先英伟达、华为、AMD 以及谷歌等芯片公司，全球更多厂商加入超节点赛道的竞争，包括微软、Meta、Amazon、中国移动、阿里巴巴、腾讯、百度、中科曙光、中兴通讯、浪潮信息、紫光股份（新华三）、沐曦股份、恒为科技等。

我们认为，全球超节点竞争格局尚未确立。英伟达目前处于领先地位，建议关注英伟达超节点供应链，包括 PCB 背板、高速铜缆、光模块、供电与液冷系统等。其次，中国厂商在超节点与 Scale up 网络领域的参与度很高，或有国内厂商在超节点领域取得领先优势，建议关注发布国产超节点的云厂商、通信设备厂商与芯片厂商。最后，基于交换机及芯片是 Scale up 网络互联的关键设备，建议关注国内交换机供应商以及交换机芯片研发商。

风险提示：(1) LLM 训练与推理技术路径变化；(2) 超节点互联方案不确定性；(3) 英伟达超节点出货量低于预期；(4) AI 应用端增长不及预期。

目 录

1. LLM 训练要求高带宽与延迟，驱动超节点成为 AI 算力网络创新方向.....	5
2. 英伟达：超节点领先优势建立在 NVLink 和 NVLink Switch	10
2.1 Scale up 网络核心技术：NVLink 与 NVLink 交换机	10
2.2 GB200 NVL72 超节点：铜缆互联，总交换容量 129.6TB/s	13
2.3 VR200 NVL72 超节点：延续 GB200 NVL72 工程技艺，总交换容量翻倍	19
2.4 总结：处于领先优势，互联 GPU 数将从 72 颗进一步向 576 颗发展	22
3. 投资建议	25
4. 风险提示	25

插图目录

图 1：Scale up 网络（左）与 Scale out 网络（右）特点对比	6
图 2：英伟达 NVL72 超节点示意	7
图 3：全球主流算力方案对应 Scale Up 协议	8
图 4：全球主流算力芯片厂商旗下 Scale up 协议特点	9
图 5：NVLink 技术规格参数对比	10
图 6：NVLink 交换机规格参数对比	10
图 7：NVLink 网状拓扑结构提供高速双向带宽	11
图 8：NVLink 交换网络的演进过程（1）	12
图 9：NVLink 交换网络的演进过程（2）	12
图 10：英伟达 GB300 NVL72 超节点外观	13
图 11：GB 200 NVL72 机柜外观与内部构件细节	14
图 12：GB200 NVL72 中计算托盘	15
图 13：GB 200 NVL72 中 NVLink 交换机托盘	15
图 14：GB 200 NVL72 中 NVLink 电缆盒	16
图 15：B200 端口 Port 示意图	17
图 16：NVLink Switch5 芯片 Port 示意图	17
图 17：GB200/300 NVL72 单层计算托架的互联拓扑	17
图 18：英伟达 GB200 NVL72 机柜后置铜线背板	18
图 19：VR200 NVL72 机柜计算托盘	19
图 20：Vera Rubin NVL72 机柜交换机托盘	19
图 21：英伟达 Vera Rubin GPU 芯片互联方式	20
图 22：VR200 NVL72 机柜中 GPU 互联拓扑结构	20
图 23：Vera Rubin NVL72 机柜交换机托盘无缆线设计	21
图 24：英伟达 Rubin NVL576 新一代 Kyber 机架	24
图 25：英伟达算力芯片发布时间表	24

表格目录

表 1: AI 大语言模型训练中多种并行计算方式对比	5
表 2: GB200 NVL72 超节点算力与通信性能	14
表 3: 英伟达超节点 Scale up 迭代路线	22

我们计划推出超节点与 Scale up 网络行业系列，详细研究英伟达、谷歌、AMD 以及华为四家头部 AI 算力芯片厂商在此领域的布局进展以及各自优势。本篇报告作为专题一，首先深度研究英伟达超节点与 Scale up 网络的进展和优势。

1. LLM 训练要求高带宽与延迟，驱动超节点成为 AI 算力网络创新方向

大语言模型 (LLM) 参数规模从千亿级向万亿级乃至十万亿级演进，跨服务器张量并行 (TP) 成为必然选择；此外混合专家 (MoE) 模型在 Transformer 架构 LLM 中的规模化应用，更使跨服务器专家并行 (EP) 成为分布式训练和推理的关键技术需求。为应对 TP 和 EP 对网络带宽与延迟的极为严苛的要求，构建超高带宽、超低延迟的 Scale up 网络（纵向扩张网络）成为业界主流技术路径。

表1：AI 大语言模型训练中多种并行计算方式对比

并行方式	带宽要求	延迟要求	说明
张量并行(TP)	数百至数千 GB/s 级	延迟要求极高	将单个运算（如矩阵乘法）拆分到不同 GPU 上运行，通常在机内完成
专家并行(EP)	数百至数千 GB/s 级	延迟要求极高	基于不同的任务选择不同专家进行训练，引入 All to All 流量，适合机内完成
流水线并行 (PP)	MB/s 至 GB/s 级	延迟要求较高	将模型的不同层划分为若干个阶段，每个阶段可以在不同的 GPU 上执行，通常在机间完成
数据并行 (DP)	GB/s 级	延迟要求较高	将同一批数据分割成多个子集，并将每个子集分配给不同 GPU 上（模型实例相同）运行，通常在机间完成

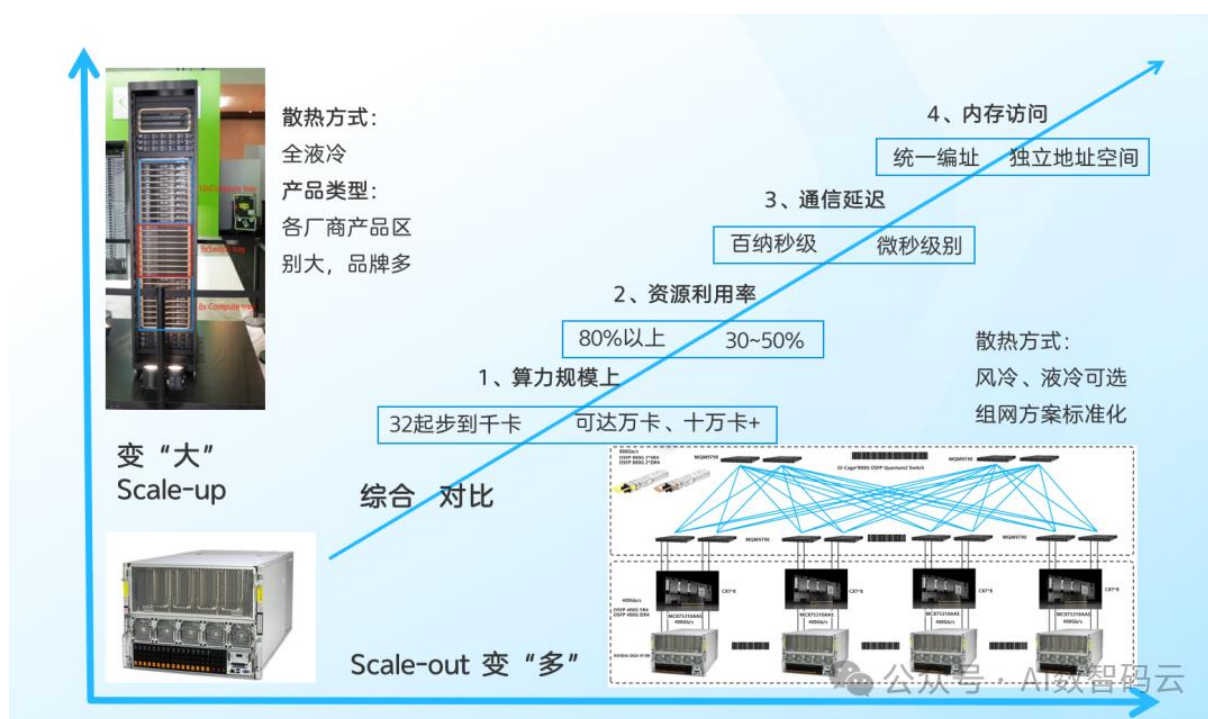
资料来源：网络技术趋势洞察公众号，东兴证券研究所

根据阿里云给出的定义为：**Scale up** 是在一定范围内，于成本和互联技术约束下实现的超高带宽互联。其范围固定且带宽是 Scale out 的数倍以上，可在协议层面优化以支持内存语义。我们对 Scale up 网络与 Scale out 网络特点对比如下：

Scale up（左）vs Scale out（右）

- 算力规模：数十卡至千卡级 vs 万卡至十万卡级；
- 资源利用率：80%以上 vs 30%-50%；
- 通信延迟：百纳秒级 vs 微秒级；
- 内存访问：统一内存或全局地址空间 vs 独立内存空间；
- 标准化：定制化程度高 vs 基于开放网络标准，相对统一。

图1：Scale up 网络（左）与 Scale out 网络（右）特点对比



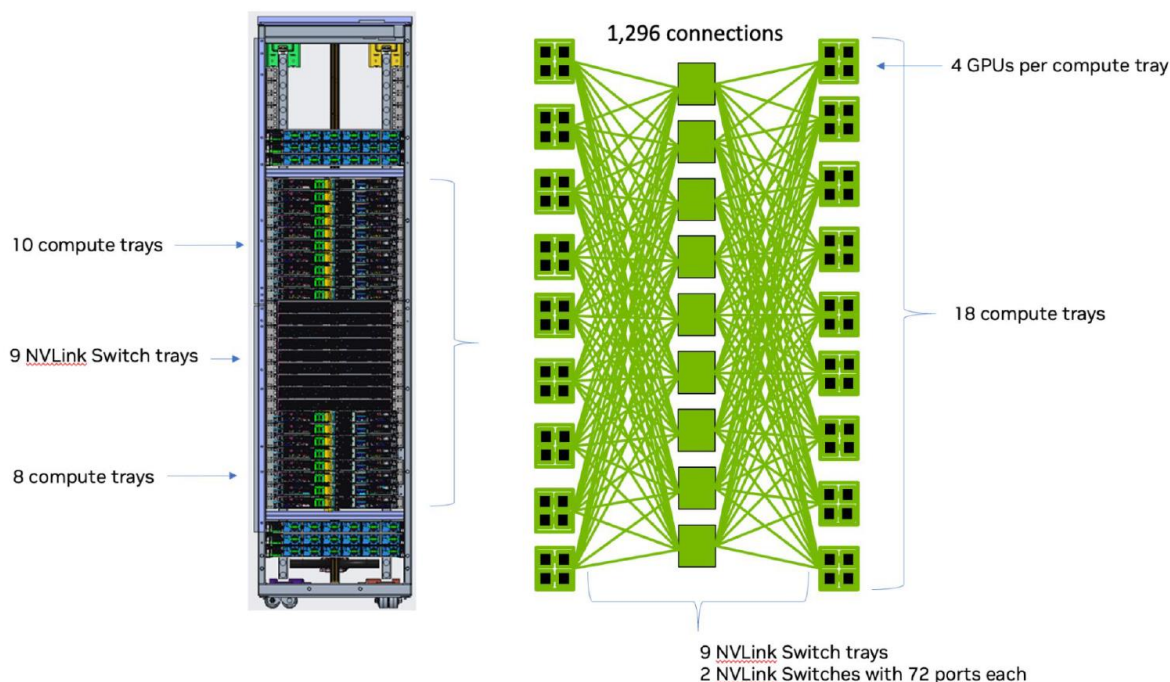
资料来源：AI 数智码云公众号，东兴证券研究所

超节点主要由计算节点、交换节点和 Scale-up 网络互联构成。通过 Scale up 网络，可将几十、上百甚至上千张 XPU 高速互联构建为超节点（SuperPoD），像一台超级 XPU 服务器一样实现高效的计算和通信协同能力。

其中 Scale up 网络互联是超节点的核心要素。Scale up 网络互联方案直接影响超节点系统的功耗、散热、成本、规模、可靠性和可维护性等关键指标。目前主流的互联方案有铜缆互联和光纤互联两大类：

- 铜缆互联方案（如英伟达的 NVL72 超节点及 NVSwitch Scale-Up 网络采用的 DAC 即无源铜缆技术）具有功耗低、成本低、可靠性高的明显优势。不过，受限于铜缆的信号传输距离，单个超节点的规模较小，目前商用的英伟达 NVL72 超节点最大支持 72 张 XPU 卡。
- 光纤互联方案（如华为的 CloudMatrix384 超节点及 Unified Bus (UB)Scale-Up 网络采用的 AOC 技术）则突破铜缆距离限制，超节点规模可以做的更大，目前商用的华为 CloudMatrix384 超节点可支持多达 384 张 XPU 卡，但这种互联技术方案也存在明显短板，如光模块功耗大，成本高，故障率高。

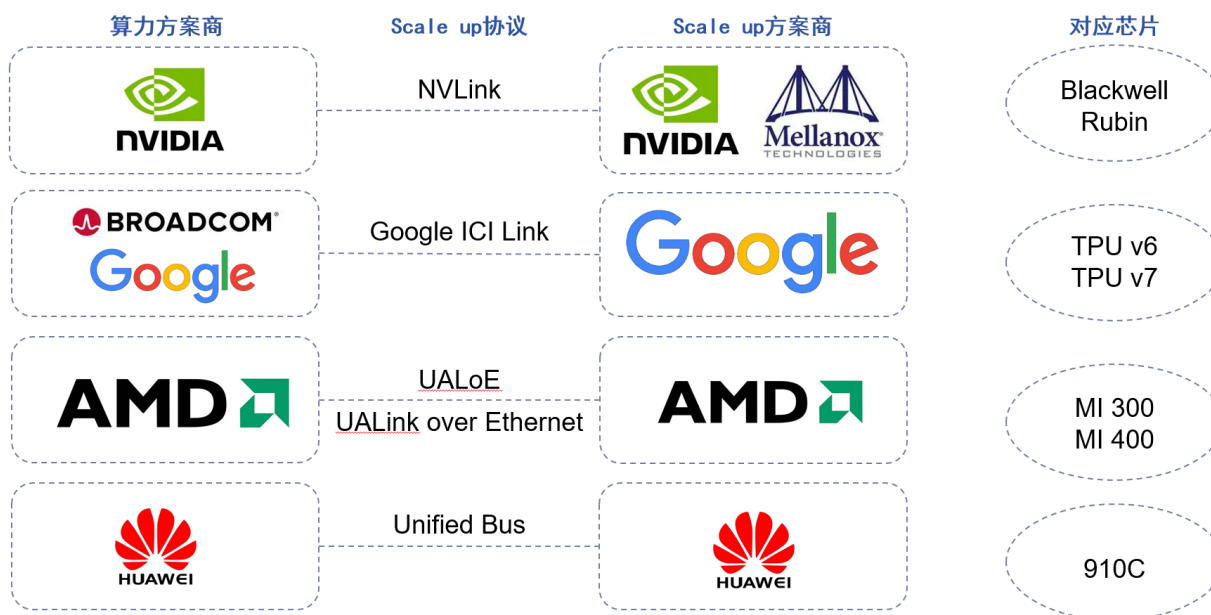
图2：英伟达 NVL72 超节点示意



资料来源：中国移动《超节点 Scale-Up 网络互联技术白皮书》，东兴证券研究所

目前英伟达、谷歌、AMD 以及华为四家头部 AI 算力芯片厂商均推出各自的 Scale up 协议。英伟达在 AI 数据中心的 Scale up 网络中采用自研的 NVLink 高速互连技术；AMD 与 AWS、思科、谷歌等公司组成超以太网联盟（UALink）；Google 采用私有 ICI 协议，机柜之间运用 OCS 光交换技术；华为推出自研的灵衢协议技术（UB）。

图3：全球主流算力方案对应 Scale Up 协议



资料来源：傅里叶的猫公众号，东兴证券研究所

Scale up 网络主要有两个技术方向。一是封闭的私有技术方向，以英伟达、Google 为典型代表，二者均采用专有协议：NVLink 仅向第三方半开放 CPU/Chiplet 接入权限；Google ICI Link 则服务于自研 TPU 集群；二是基于 Ethernet 的开放技术方向，以各大互联网和云计算公司以及一些 GPU 芯片公司为代表。开放标准以 UALink 和 华为灵衢 为代表，UALink 基于标准以太网组件打造开放互联协议，华为灵衢从 2.0 版本起转向开放标准。目前两者均处于生态建设初期。

图4：全球主流算力芯片厂商旗下 Scale up 协议特点

	NVIDIA NVLink™	Google ICI Link	ULTRA ACCELERATOR LINK™	灵衢 UnifiedBus
主导方	英伟达	谷歌	超以太网联盟（Meta、微软、英特尔、AMD 等）	华为
协议性质	专有协议	专有协议	开放标准	自灵衢 2.0 起开放标准
技术标准	基于 SerDes 的私有协议	ICI 协议（支持 3D 环面拓扑）	使用标准以太网组件	基于高速 SerDes 的物理层
连接形式	机柜内通过铜缆实现 GPU 互联	结合铜缆与光缆连接，立方体内部使用铜缆连接，立方体外部使用光缆连接	同时支持铜缆和光缆	机柜内使用电缆互联，不采用光模块
生态支持	NVLink-Fusion，通过半定制需通过 XLA 编译器将计算图 CPU 或 Chiplet 向第三方 GPU 开放生态系统	转为 TPU 机器码，深度绑定 Google 生态	兼容多厂商 GPU（如 AMD Instinct MI 系列、英特尔 Gaudi 等）	全栈开放协议，处于建设初期，第三方商用 GPU 产品尚未大规模上市

资料来源：Semi Analysis，CSDN，东兴证券研究所

2. 英伟达：超节点领先优势建立在 NVLink 和 NVLink Switch

2.1 Scale up 网络核心技术：NVLink 与 NVLink 交换机

NVLink 与 NVLink 交换机是英伟达构建单机柜 Scale up 网络的核心技术组合。二者协同演进，从早期点对点互联发展到如今全互联通信，并支持多代 GPU 架构算力芯片。2026 年 1 月，英伟达发布第六代 NVLink 以及 NVLink 交换机，两者支持最新的 Rubin 架构。从性能指标看，第六代 NVLink 交换机支持的 GPU-to-GPU 通信带宽为 3.6TB/s；在 VR NVL72 系统中提供 260TB/s 聚合带宽。其中每 GPU 的 NVLink 带宽保持不变，与 NVLink5.0 一致，仍为 100GB/s。

图5：NVLink 技术规格参数对比

	NVLink NVLink 交换机		
	第四代	第五代	第六代
每 GPU 的 NVLink 带宽	900GB/s	1,800GB/s	3,600 GB/s
每 GPU 的最大链路数	18	18	36
支持的 NVIDIA 架构	NVIDIA Hopper™ 架构	NVIDIA Blackwell 架构	NVIDIA Rubin 平台

资料来源：英伟达官网，东兴证券研究所

图6：NVLink 交换机规格参数对比

	NVLink NVLink 交换机		
	NVLink 4 交换机	NVLink 5 交换机	NVLink 6 交换机
NVLink GPU 域	8	8 72	8 72
NVLink 交换机 GPU 到 GPU 带宽	900 GB/s	1,800 GB/s	3,600 GB/s
总聚合带宽	7.2 TB/s	130 TB/s (NVL72)	260 TB/s (NVL72)
支持的 NVIDIA 架构	NVIDIA Hopper™ 架构	NVIDIA Blackwell 架构	NVIDIA Rubin 平台

资料来源：英伟达官网，东兴证券研究所

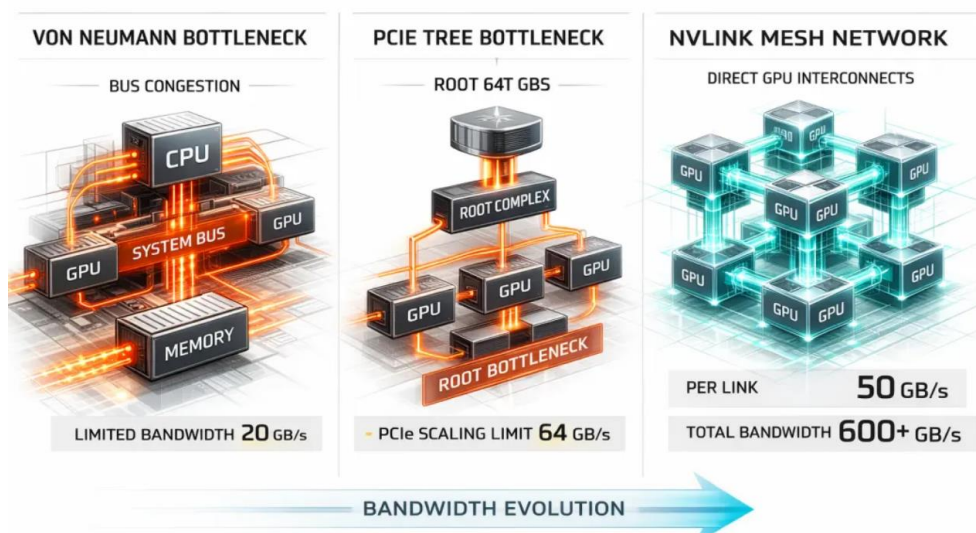
NVLink 重新设计通信架构，推出网状拓扑理念。为实现 AI 训练集群高带宽与低延迟数据传输，NVLink 允许 GPU 之间形成多对多的直接通信网络，每个 GPU 都可以同时与多个其他 GPU 建立高速通信链路。NVLink 协议创新如下：

在物理层面，NVLink 采用差分信号传输技术，具有高带宽和高抗干扰性能。每个链路由多对差分信号线组成，每对信号线负责传输一个方向的数据。SerDes 模块是 NVLink 物理层的核心组件，负责将并行数据转换为高速串行流，并在接收端进行反向转换。NVLink 的 SerDes 设计采用时钟数据恢复技术，以及集成复杂的自适应均衡电路。

在链路层，NVLink 定义多种类型的符号，包括数据符号、控制符号和填充符号，实现复杂的通信协议功能；设计精细的信用机制，实现不同优先级的流量调度。

除此之外，NVLink 其他创新之处包括多 Lane 绑定技术、统一内存空间等。

图7：NVLink 网状拓扑结构提供高速双向带宽



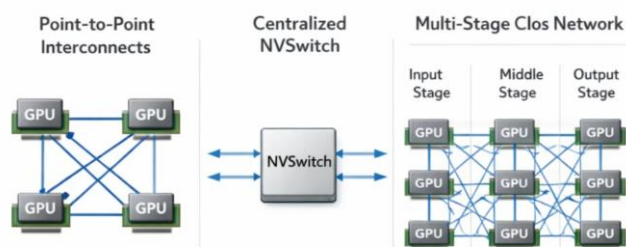
资料来源：仰望 7866 公众号，东兴证券研究所

NVSwitch 是实现 Scale up 网络复杂交换的关键设备。

早期的 NVLink 实现主要采用点对点连接模式，GPU 之间通过直接的串行链路进行通信。当系统包含多个 GPU 时，点对点模式的连接复杂度呈平方级增长。

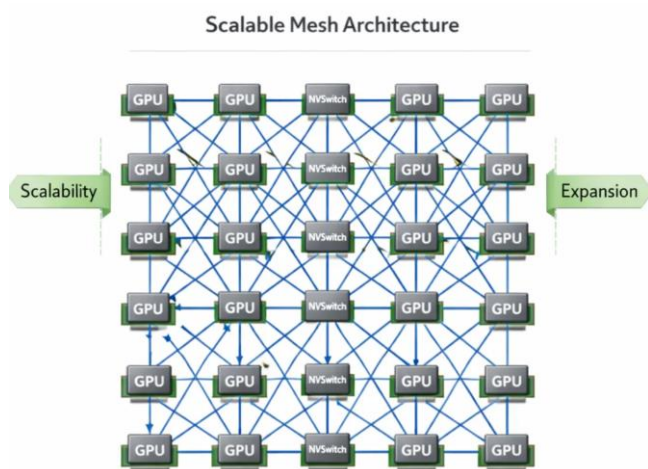
作为专门的交换芯片，NVSwitch 可以提供多端口的高速交换能力。NVLink 的交换网络采用多阶 Clos 网络架构，Clos 网络通过多级交换结构实现输入端口到输出端口的任意连接。

图8：NVLink 交换网络的演进过程（1）



资料来源：仰望 7866 公众号，东兴证券研究所

图9：NVLink 交换网络的演进过程（2）



资料来源：仰望 7866 公众号，东兴证券研究所

2.2 GB200 NVL72 超节点：铜缆互联，总交换容量 129.6TB/s

目前英伟达超节点已经推出成熟方案，在行业中处于领先地位。2024-2026 年，英伟达陆续推出 GH200 NVL72、GB200/ GB300 NVL72、VR200 NVL72 三代超节点。

- **Hopper 架构开启超节点 Scale up 初步探索。**GH200 通过 NVLink 和 NVLink-C2C（Chip-to-Chip）技术，使得每个 GPU 可以访问其他所有 CPU 和 GPU 芯片的内存，实现 GPU 与 CPU 内存统一编址。
- **Blackwell 架构推动 Scale up 标准化。**GB200 NVL72 将 Scale-up 规模稳定在 72 个 GPU/机柜，形成可复制标准化方案。NVL72 由 18 个 Compute Tray（计算托架）和 9 个 Switch Tray（网络交换托架）构成。其中，Compute Tray 是计算核心单元，负责提供强大的计算能力；Switch Tray 是高速通信枢纽，用于实现 GPU 之间的高速数据交换。NVL72 背板通过“NVLink5 私有协议 + 铜线缆”将 18 个 Compute Tray 中的 72 颗 B200 GPU 和 9 个 Switch Tray 中的 18 颗 NVSwitch 芯片进行满带宽全连接。
- **Rubin 架构推动 Scale up 方案带宽倍增。**2026 年 1 月 CES 展会，英伟达发布 Rubin 架构 VR200 NVL72。其中 NVLink 6 Switch 实现单 GPU 的互连带宽提升至 3.6 TB/s，上代为 1.8TB/s。Scale out 方面，Spectrum-6 交换机支持 CPO（共封装光学）技术，将 32 个 1.6Tb/s 硅光光学引擎与交换芯片直接封装集成。

图10：英伟达 GB300 NVL72 超节点外观



资料来源：传热之道公众号，东兴证券研究所

目前全球算力芯片公司进入芯片性能与超节点性能并行竞争的新阶段。GB200 NVL72 作为全球超节点发展的标杆产品，我们将从多个维度拆解其硬件构成以及重点性能指标。

从算力和通信性能看：GB200 NVL72 提供 180 PFLOP 的 TF32 Tensor Core 算力，总内存容量 13.8TB，内存带宽 576TB/s；Scale up 带宽 64800 单向 GB/s。

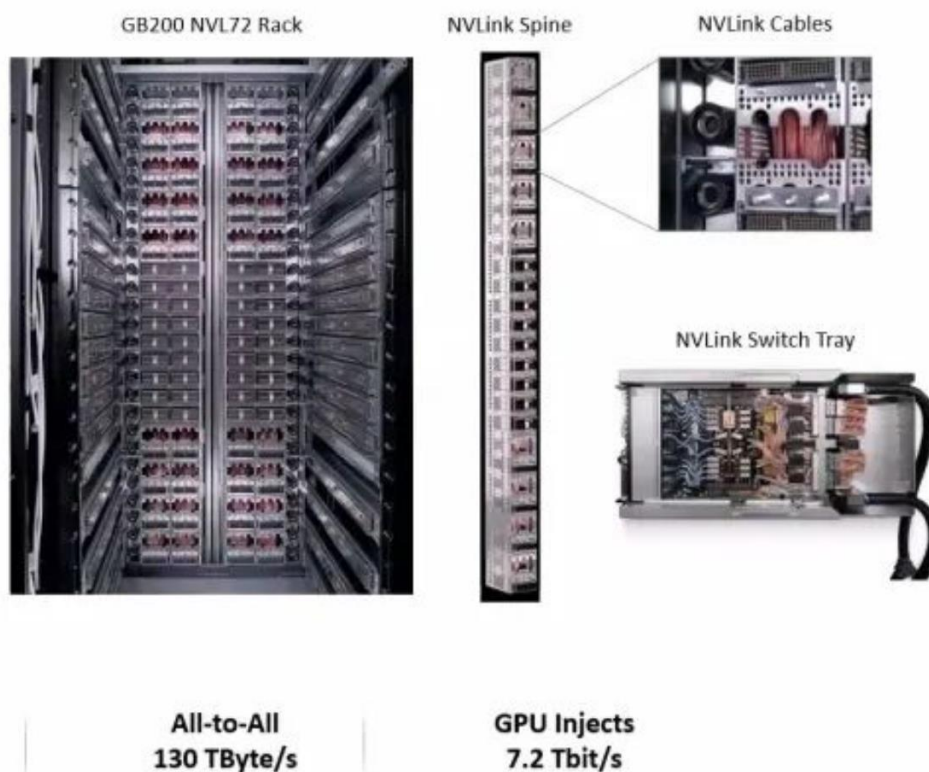
表2：GB200 NVL72 超节点算力与通信性能

	单位	GB200 NVL72
算力（TF32 Tensor 核心）	PFLOPS	180
HBM 内存	TB	13.4
HBM 带宽	TB/s	576
Scale up 带宽	单向 GB/s	64800
Scale up 计算单元	GPUs	72
功耗	KW	145

资料来源：SemiAnalysis, Nvidia, 华为, 东兴证券研究所

除了算力与通信性能，尺寸、重量、功耗均是超节点 TCO（总体拥有成本）的关键影响因素。GB200 NVL72 机柜尺寸为长 1068 毫米、宽 600 毫米、高 2495 毫米；重约 1.36 吨；功耗 145KW。

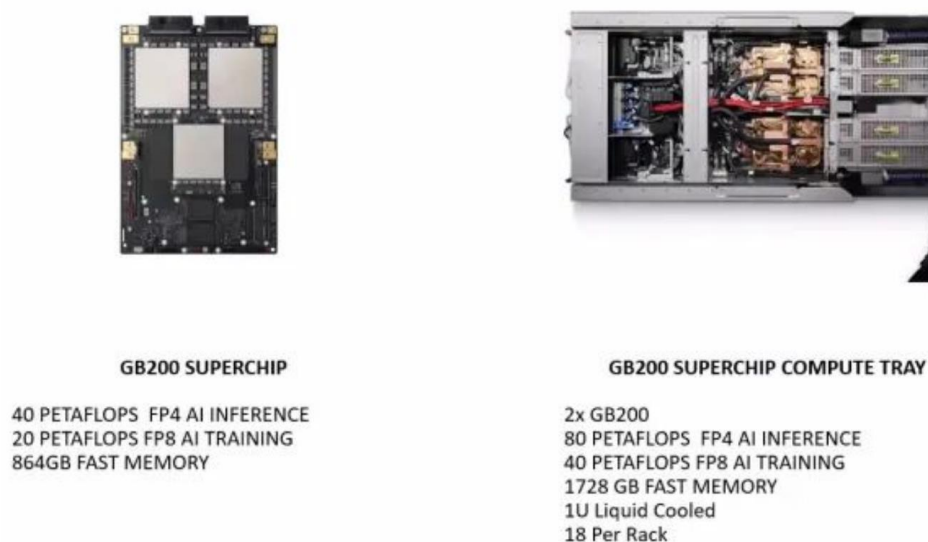
图11：GB 200 NVL72 机柜外观与内部构件细节



资料来源：光芯公众号，东兴证券研究所

单台 GB 200 NVL72 机柜有 18 个计算节点。GB200 NVL72 超节点主要由 18 个 Compute Tray(计算托盘)和 9 个 Switch Tray(网络交换托盘)构成。每个计算托盘容纳 4 颗 B200 GPU 和 2 颗 Grace CPU，构成两个 GB200 超级芯片。

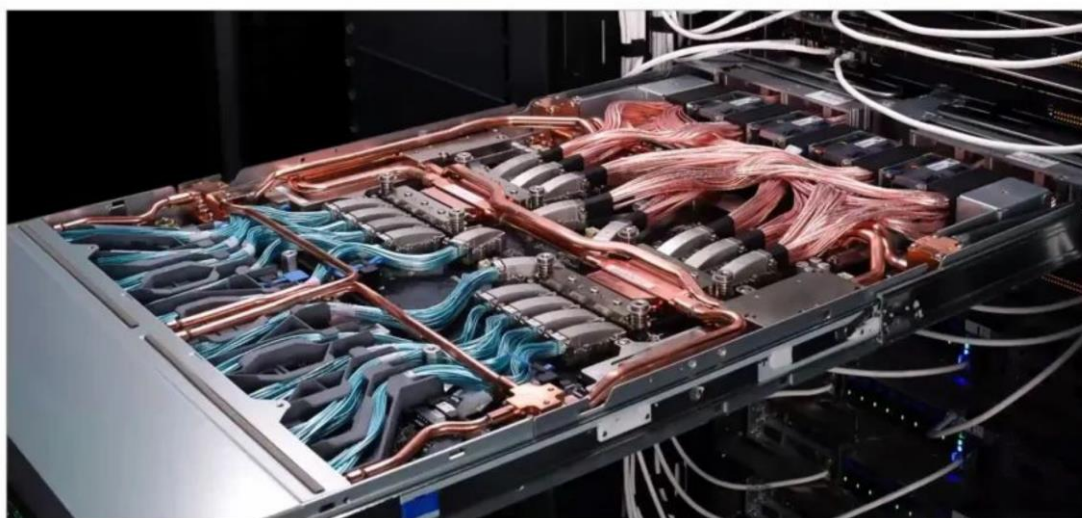
图12：GB200 NVL72 中计算托盘



资料来源：光芯公众号，东兴证券研究所

GB 200 NVL72 机柜有 9 个网络交换托盘。每个网络交换托盘包含两颗 NVLINK Switch5 芯片，合计 18 颗 NVSwitch5 芯片。单颗 NVSwitch5 芯片交换容量为 7.2TB/s，总交换容量 129.6TB/s。网络交换托架中金色电缆用于 NVLink 连接，与电缆盒相连，机箱前面的蓝色电缆用于 OSFP 接口，实现不同版本的扩展。

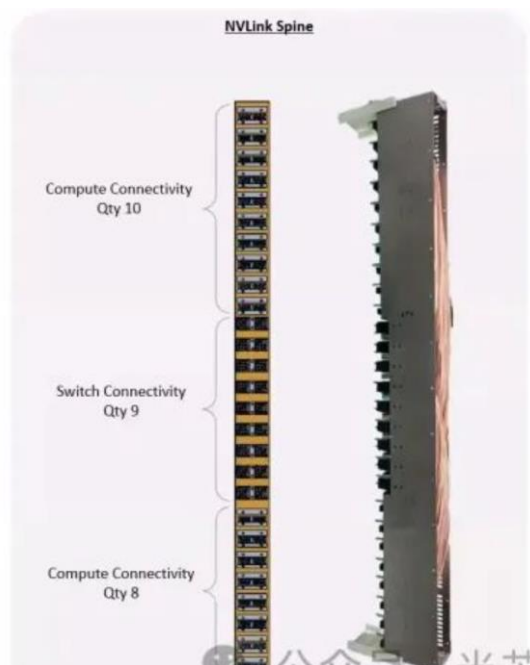
图13：GB 200 NVL72 中 NVLink 交换机托盘



资料来源：光芯公众号，东兴证券研究所

电缆盒负责垂直方向信号重组。电缆盒有 8 个底部连接器和 10 个顶部连接器，每个连接器可处理一个 GPU 的全部带宽。

图14：GB 200 NVL72 中 NVLink 电缆盒



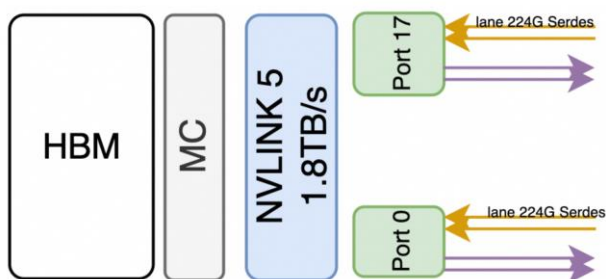
资料来源：光芯公众号，东兴证券研究所

GB200 NVL72 实现 72 颗 B200 完全互联，总交换带宽 129.6TB/s。

计算节点访存带宽为 7.2TB/s：B200 设置 18 个端口（Port）。每个端口采用 224G Serdes，由四对差分线构成。每个端口的传输速率为 $200\text{Gbps} \times 4$ （4 对差分线）/8 = 100GB/s（双向）。每个计算托盘容纳 4 颗 B200 GPU，则每个计算节点 72 个 NVLink5 Port，总访存带宽为 7.2TB/s。

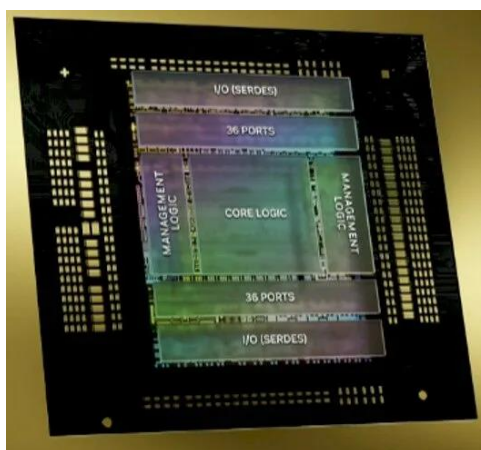
交换节点访存带宽为 14.4TB/s：NVSwitch5 芯片由 72 个 NVLINK Port（上下各 36 个 Port）。同样，每个 Port 采用双路 200Gbps 速率的 SerDes 高速串行接口，则每个 Port 带宽为 100GB/s。每个交换托盘两颗 NVLINK Switch5 芯片。每个交换节点 144 个 NVLINK Port，总访存带宽为 14.4TB/s。

图15：B200 端口 Port 示意图



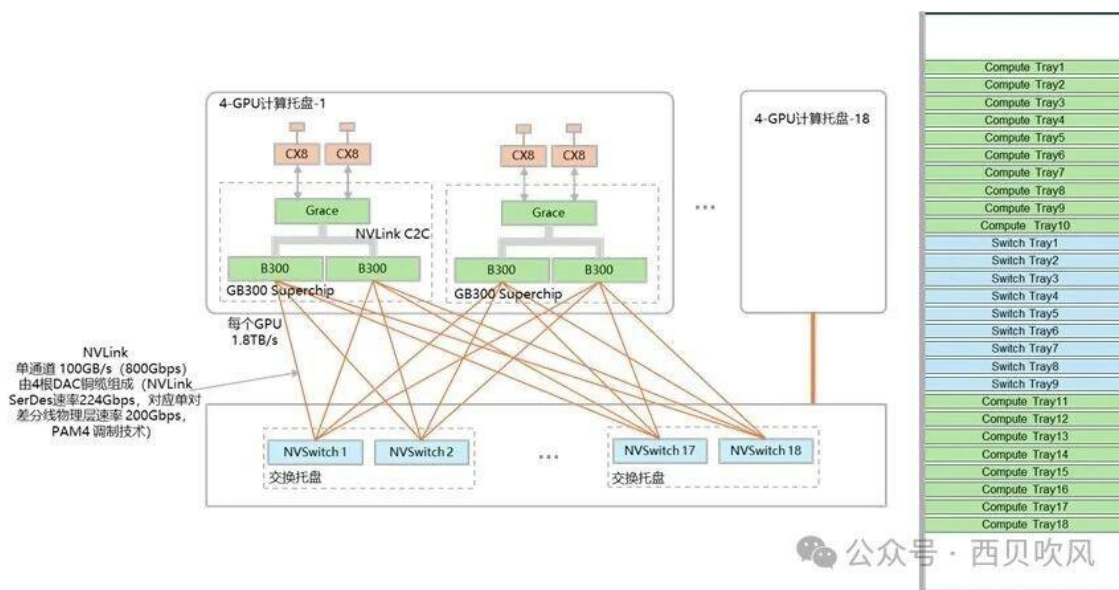
资料来源：zartbot 公众号，东兴证券研究所

图16：NVLINK Switch5 芯片 Port 示意图



资料来源：zartbot 公众号，东兴证券研究所

图17：GB200/300 NVL72 单层计算托架的互联拓扑



资料来源：西贝吹风公众号，东兴证券研究所

GB200 NVL 72 Scale up 方案中以铜缆互联为主。GB200 NVL72 在互联方案中主要采用直连铜缆 (DAC)，在某些特殊场景（如跨托盘连接或需要稍长传输距离的场景）中，会采用 ACC 铜缆。ACC（主动铜缆，在 DAC 基础上增加有源信号处理芯片）的信号增强能力可以弥补 DAC 在较长距离传输时的信号衰减问题，确保数据传输的稳定性和可靠性。

在 GB200 NVL 72 中所需铜缆数量： $18 \text{ (托盘数量)} \times 4 \text{ (GPU 数量)} \times 4 \text{ (GPU 到 NVSwitch 单端口铜缆数量)} \times 18 \text{ (NVSwitch 数量)} = 5184 \text{ 根}$ 。（100GB/s 单端口由 4 根 DAC 铜缆组成）

图18：英伟达 GB200 NVL72 机柜后置铜线背板



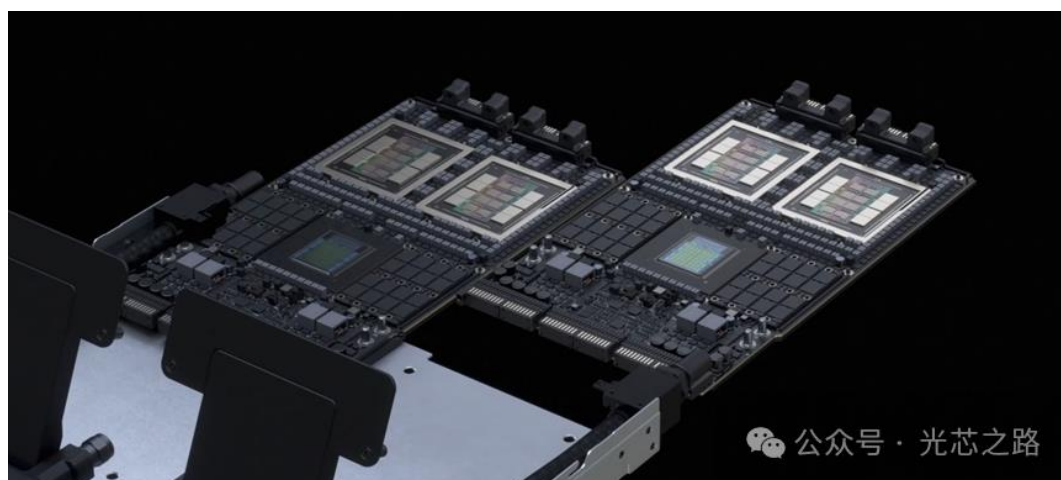
资料来源：英伟达 GTC，东兴证券研究所

2.3 VR200 NVL72 超节点：延续 GB200 NVL72 工程技艺，总交换容量翻倍

2026 年 1 月 6 日在 CES 2026 展会上，英伟达发布新一代超节点 VR200 NVL72。我们认为，相比 GB200 NVL72，新一代 VR200 NVL72 属于连续性创新，而非破坏性创新。具体对比如下：

计算节点：Rubin NVL72 机架通过无缆互联架构整合 18 个计算托盘，每个托盘 2 颗超级芯片，每颗超级芯片集成 1 个 Vera CPU 与 2 块 Rubin GPU，共 72 GPU 与 36 CPU。

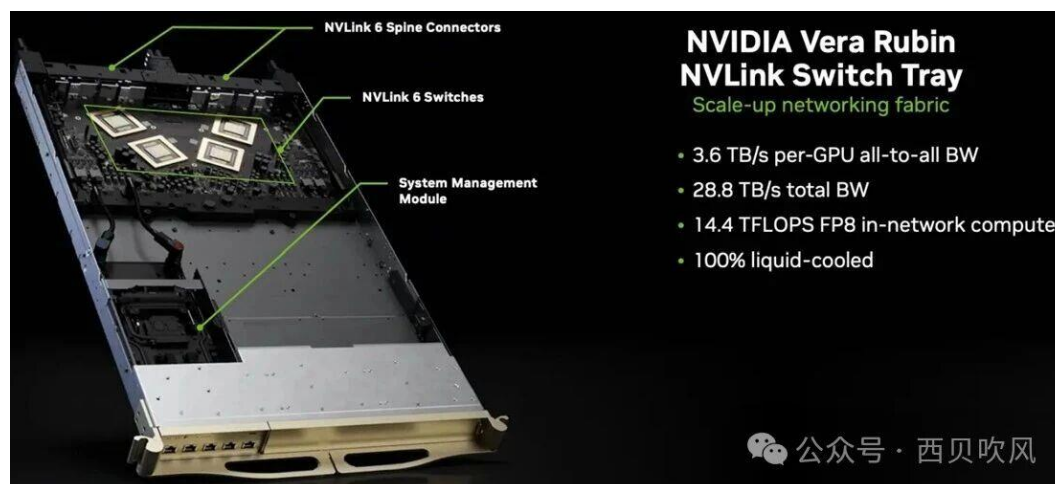
图19：VR200 NVL72 机柜计算托盘



资料来源：光芯之路公众号，东兴证券研究所

交换节点：VR200 NVL72 配置 9 个交换托盘，每个托盘集成 4 颗第六代 NVSwitch 芯片，全机柜部署 36 颗 NVSwitch。相比 GB200 NVL72，NVSwitch 芯片数量实现翻倍。单颗第六代 NVSwitch 交换容量为 7.2TB/s，相比 NVSwitch5 芯片，保持不变。

图20：Vera Rubin NVL72 机柜交换机托盘



资料来源：西北吹雪公众号，东兴证券研究所

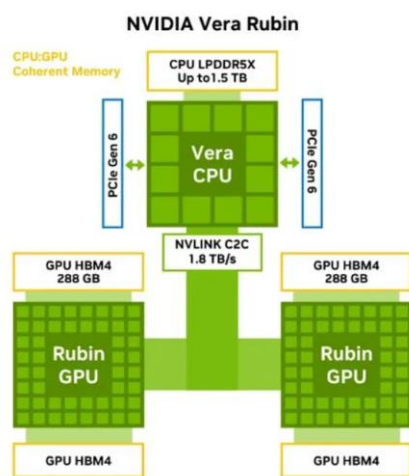
VR200 NVL72 Scale up 方案实现总交换容量 259.2TB/s，对比 GB200 NVL72，提升一倍。

计算节点：VR200 设置 72 个端口。每个端口带宽 100GB/s。每个计算托盘 2 颗 VR200 GPU，则每个计算节点 144 个 NVLink 6.0 端口，总访存带宽为 14.4TB/s。

交换节点：NVSwitch6 芯片 72 个 NVLink 6.0 端口。每个 NVLinkPort 交换容量 100GB/s。每个交换托盘 4 个 NVLink 6 Switch 芯片，每个交换节点 288 个 NVLink Port，总访存带宽为 28.8TB/s。

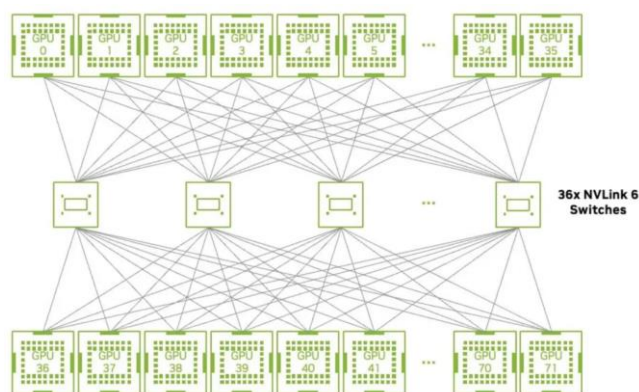
此外，VR200 NVL72 依托 NVLink-C2C 实现 1.8TB/s CPU-GPU 互联，相比 GB200 NVL72 的 NVLink-C2C 的速率为 900GB/s，提升一倍。

图21：英伟达 Vera Rubin GPU 芯片互联方式



资料来源：英伟达官网，东兴证券研究所

图22：VR200 NVL72 机柜中 GPU 互联拓扑结构



资料来源：英伟达官网，东兴证券研究所

VR200 NVL72 Scale up 方案延续铜缆互联方案。稍有不同之处，VR200 用中板取代计算托盘内部的线缆，中板采用覆铜板技术。此外，基于 Rubin 平台 NVLink6.0 升级至 448G SerDes 通道速率。因此，GPU 到每个 NVSwitch 铜缆连接由 4 根变为 2 根。

主干铜缆数量 = 18（托盘数量）* 4（GPU 数量）* 2（GPU 到 NVSwitch 铜缆数量）* 36（NVSwitch 数量）
= 5184 根。

图23：Vera Rubin NVL72 机柜交换机托盘无缆线设计



资料来源：西北吹雪公众号，东兴证券研究所

2.4 总结：处于领先优势，互联 GPU 数将从 72 颗进一步向 576 颗发展

在超节点方案上，英伟达处于领先优势。2024-2025 年，英伟达陆续推出 GH200 NVL72、GB200/ GB300 NVL72 等成熟超节点解决方案。根据大摩预测，2025 年英伟达 GB200/300 NVL72 出货量约 2800 台。展望 2026-2027 年，英伟达计划推出 Vera Rubin NVL144 和 Rubin Ultra NVL576。互联 GPU 数将从 72 颗进一步向 576 颗发展。届时，英伟达将在新一代 Kyber 机架架构中引入 NVLink Switch Blade（NVLink 交换机刀片），通过 PCB 中板替代传统 5000+ 根有源铜缆。可以看到，Rubin Ultra NVL576 仍具有较强的工程创新能力。

英伟达超节点的优势建立在 NVLink 和 NVLink Switch。为实现 AI 训练集群高带宽与低延迟数据传输，NVLink 重新设计通信架构，并引入一系列先进技术，包括网状拓扑、差分信号传输、流量调度信用机制、多 Lane 绑定技术、统一内存空间等。截止 2025 年，NVLink 5 Switch 实现支持单 GPU 到 GPU 带宽 1800GB/s，可构建 72 GPU 的 NVLink 域，总带宽达 130 TB/s（双向），支持 72 GPU 全互联通信。在后续计划中，NVSwitch Gen6 和 Gen7 的 GPU-to-GPU 通信带宽继续升级为 3.6TB/s。

但另一方面，Scale up 网络兴起源于满足大模型分布式训练和推理中的张量并行(TP)与专家并行(EP)。目前 AI 产业也在探索降低 TP 与 EP 规模的技术方案，从而降低 Scale up 网络规模的上限。我们认为，Scale up 网络的发展空间或限制英伟达在超节点领域的领先优势。为保持领先优势，实现 Scale up 网络和 Scale out 网络融合或将成为英伟达超节点新的发展趋势。

表3：英伟达超节点 Scale up 迭代路线

架构	Blackwell Ultra	Vera Rubin NVL72	Vera Rubin NVL144	Rubin Ultra NVL576	Feunman
首发时间	2025-03	2026-01	预计 2026 下半年	预计 2027 年	预计 2028 年
核心平台	GB300 NVL72	VR200 NVL72	VR200 NVL144	Rubin Ultra NVL576	Feynman NVL1152
计算托盘	18 个（单盘 4GPU+2CPU）	18 个（单盘 4GPU+2CPU）	36 个（单盘 4GPU+2CPU）	72 个（单盘 8GPU+4CPU）	144 个（单盘 8GPU+4CPU）
CPU	36 Grace CPUs（72 核）	36 Vera CPUs（72 核）	72 Vera CPUs（88 核）	288 Vera Ultra CPUs （176 核）	576 Feynman CPUs （256 核）
	单颗内存带宽 3.6Tpbs	单颗内存带宽 4.8Tpbs	单颗内存带宽 4.8Tpbs	单颗内存带宽 9.6Tpbs	单颗内存带宽 19.2Tpbs
	NVLink C2C 0.9Tpbs	NVLink C2C 1.8Tpbs	NVLink C2C 1.8Tpbs	NVLink C2C 3.6Tpbs	NVLink C2C 7.2Tpbs
GPU	72 GB300 GPUs	72 VR200 GPUs	144 VR200 GPUs	576 VR300 GPUs	1152 Feynman GPUs
	单颗 288G HBM3E	单颗 512G HBM3E	单颗 512G HBM3E	单颗 1TB HBM4E	单颗 2TB HBM5E
	单颗 MVFP4 15PFLOPS	单颗 MVFP4 50PFLOPS	单颗 MVFP4 50PFLOPS	单颗 MVFP4 100PFLOPS	单颗 MVFP4 200PFLOPS
	铜缆背板	铜缆背板	铜缆背板+板载无源光引擎（非 CPO）	3.2T CPO 硅光（规划）	6.4T CPO 硅光（规划）
Scale up	18 个 NVLink5	36 个 NVLink6	72 个 NVLink6	144 个 NVLink7	72 个 NVLink8 硅光交换机
	单个 144*200G 28.8T	单个 72*400G 28.8T	单个 72*400G 28.8T	单个 144*800G 115.2T	单个 288*1.6T 460.8T
	GPU 侧 NVLink 带宽 18*1.8TBps	GPU 侧 NVLink 带宽 18*3.6TBps	GPU 侧 NVLink 带宽 18*3.6TBps	GPU 侧 NVLink 带宽 36*7.2TBps	GPU 侧 NVLink 带宽 72*14.4TBps

架构	Blackwell Ultra	Vera Rubin NVL72	Vera Rubin NVL144	Rubin Ultra NVL576	Feunman
Scale out	Spectrum-5 800G OSFP	Spectrum-6 CPO 硅光	Spectrum-6 CPO 硅光	Spectrum-7 CPO 硅光	Spectrum-8 CPO 硅光
	可插拔				
	64*800G 51.2T	128*800G 102.4T	128*800G 102.4T	256*1.6T 409.6T	512*3.2T 1638.4T

注：28.8Tbps=3.6TBps

信息来源：光芯之路公众号，东兴证券研究所

Rubin NVL576 由单个计算柜配置一个 Kyber SideCar 机柜构成。计算柜内由 4 个 NVL144 的计算机框构成，每个计算框包含 18 个 ComputeTray；Switch Blade 垂直插入机架后部，与计算刀片通过中板直接连接。

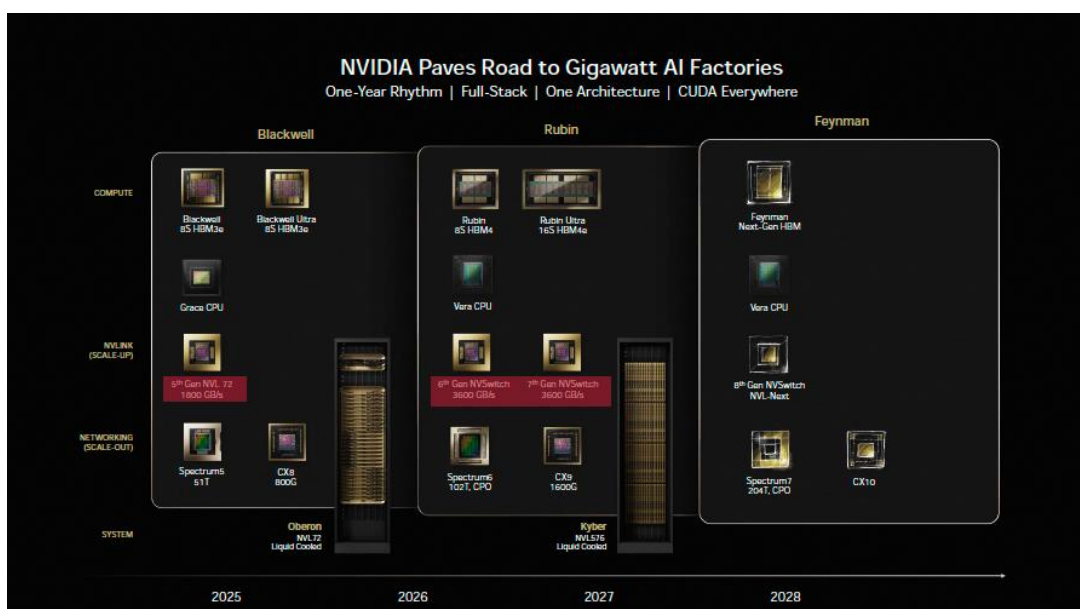
图24：英伟达 Rubin NVL576 新一代 Kyber 机架



资料来源：GTC2025，东兴证券研究所

后续 NVSwitch Gen6 和 Gen7 的 GPU-to-GPU 通信带宽为 3.6TB/s。

图25：英伟达算力芯片发布时间表



资料来源：GTC2025，东兴证券研究所

3. 投资建议

自 2025 年开始，超节点成为 AI 算力网络重要的技术创新方向。从 AI 基建竞争维度，AI 芯片厂商从芯片算力性能竞争延续至芯片+Scale up 网络的双战场。因此，除了原先英伟达、华为、AMD 以及谷歌等芯片公司，全球更多厂商加入超节点赛道的竞争，包括微软、Meta、Amazon、中国移动、阿里巴巴、腾讯、百度、中科曙光、中兴通讯、浪潮信息、紫光股份（新华三）、沐曦股份、恒为科技等。

我们认为，全球超节点竞争格局尚未确立。英伟达目前处于领先地位，建议关注英伟达超节点供应链，包括 PCB 背板、高速铜缆、光模块、供电与液冷系统等。其次，中国厂商在超节点与 Scale up 网络领域的参与度很高，或有国内厂商在超节点领域取得领先优势，建议关注发布国产超节点的云厂商、通信设备厂商与芯片厂商。最后，基于交换机及芯片是 Scale up 网络互联的关键设备，建议关注国内交换机供应商以及交换机芯片研发商。

4. 风险提示

- （1）LLM 训练与推理技术路径变化；
- （2）超节点互联方案不确定性；
- （3）英伟达超节点出货量低于预期；
- （4）AI 应用端增长不及预期。

分析师简介

石伟晶

首席分析师，覆盖传媒、互联网、云计算、通信等行业。上海交通大学工学硕士。10 年证券从业经验，曾供职于华创证券、安信证券，2018 年加入东兴证券研究所。

分析师承诺

负责本研究报告全部或部分内容的每一位证券分析师，在此申明，本报告的观点、逻辑和论据均为分析师本人研究成果，引用的相关信息和文字均已注明出处。本报告依据公开的信息来源，力求清晰、准确地反映分析师本人的研究观点。本人薪酬的任何部分过去不曾与、现在不与、未来也将不会与本报告中的具体推荐或观点直接或间接相关。

风险提示

本证券研究报告所载的信息、观点、结论等内容仅供投资者决策参考。在任何情况下，本公司证券研究报告均不构成对任何机构和个人的投资建议，市场有风险，投资者在决定投资前，务必要审慎。投资者应自主作出投资决策，自行承担投资风险。

免责声明

本研究报告由东兴证券股份有限公司研究所撰写，东兴证券股份有限公司是具有合法证券投资咨询业务资格的机构。本研究报告中所引用信息均来源于公开资料，我公司对这些信息的准确性和完整性不作任何保证，也不保证所包含的信息和建议不会发生任何变更。我们已力求报告内容的客观、公正，但文中的观点、结论和建议仅供参考，报告中的信息或意见并不构成所述证券的买卖出价或征价，投资者据此做出的任何投资决策与本公司和作者无关。

我公司及报告作者在自身所知情的范围内，与本报告所评价或推荐的证券或投资标的的存在法律禁止的利害关系。在法律许可的情况下，我公司及其所属关联机构可能会持有报告中提到的公司所发行的证券头寸并进行交易，也可能为这些公司提供或者争取提供投资银行、财务顾问或者金融产品等相关服务。本报告版权仅为我公司所有，未经书面许可，任何机构和个人不得以任何形式翻版、复制和发布。如引用、刊发，需注明出处为东兴证券研究所，且不得对本报告进行有悖原意的引用、删节和修改。

本研究报告仅供东兴证券股份有限公司客户和经本公司授权刊载机构的客户使用，未经授权私自刊载研究报告的机构以及其阅读和使用者应慎重使用报告、防止被误导，本公司不承担由于非授权机构私自刊发和非授权客户使用该报告所产生的相关风险和法律责任。

行业评级体系

公司投资评级（A股市场基准为沪深300指数，香港市场基准为恒生指数，美国市场基准为标普500指数）：
以报告日后的6个月内，公司股价相对于同期市场基准指数的表现为标准定义：

强烈推荐：相对强于市场基准指数收益率15%以上；

推荐：相对强于市场基准指数收益率5%~15%之间；

中性：相对于市场基准指数收益率介于-5%~+5%之间；

回避：相对弱于市场基准指数收益率5%以上。

行业投资评级（A股市场基准为沪深300指数，香港市场基准为恒生指数，美国市场基准为标普500指数）：
以报告日后的6个月内，行业指数相对于同期市场基准指数的表现为标准定义：

看好：相对强于市场基准指数收益率5%以上；

中性：相对于市场基准指数收益率介于-5%~+5%之间；

看淡：相对弱于市场基准指数收益率5%以上。

东兴证券研究所

北京

西城区金融大街5号新盛大厦B座16层

邮编：100033

电话：010-66554070

传真：010-66554008

上海

虹口区杨树浦路248号瑞丰国际大厦23层

邮编：200082

电话：021-25102800

传真：021-25102881

深圳

福田区益田路6009号新世界中心46F

邮编：518038

电话：0755-83239601

传真：0755-23824526