

电子行业深度报告

AI 基建，光板铜电—GTC 前瞻

Serdes, Rubin Ultra&CPO 交换机详解

增持（维持）

2026 年 02 月 25 日

证券分析师 陈海进

执业证书：S0600525020001

chenhj@dwzq.com.cn

研究助理 解承堯

执业证书：S0600125020001

xiechy@dwzq.com.cn

投资要点

■ **重点关注 2026 年 M9 产业链、光互联产业链投资机遇。**本系列首篇报告《AI 基建，光板铜电—光&铜篇，主流算力芯片 Scale up&out 方案全解析》深度解析了 2026 年四款核心算力芯片的 Scale-up 与 Scale-out 组网架构，并量化测算了 AI 服务器中 PCB、高速铜缆、光模块等信号传输介质的结构配比。本篇作为系列第二篇，暨 GTC 2026 大会前瞻，我们立足英伟达 SerDes 技术演进，前瞻推演 2026-2027 年整机柜架构及 CPO 产业化进程。穿透产业链繁杂噪音，本报告纯粹从技术底层出发，系统剖析了 SerDes 迭代面临的速率瓶颈与功耗约束，从而判断出 PCB 材料向 M9 等级升级、光电共封装及“光入柜内”的确定性技术趋势。因此我们建议 2026 年重点关注 PCB M9 材料产业链投资机遇，以及 CPO、光入柜内所对应的光芯片、光器件等领域的投资机遇。

■ **SerDes 代际跃迁，驱动算力互联介质升级。**算力芯片的互联带宽已成为衡量其性能的核心指标，而 SerDes（高速串行解串器）作为高速 IO 端口的核心技术组件，其速率迭代直接决定了芯片互联带宽的上限。以英伟达为例，NVLink SerDes 已从 Ampere 架构的 56Gbps 演进至 Blackwell 架构的 224Gbps，支撑单芯片互联带宽实现代际跨越。然而，SerDes 速率的持续提升对 AI 服务器，交换机的信号传输介质提出严苛挑战。从速率瓶颈看，224G 以上信号高频衰减剧增，驱动 PCB 覆铜板向 M9 级别升级；从功耗瓶颈看，SerDes 功耗占比随速率攀升，光互联亟需通过光电近封装、共封装技术缩短与交换芯片的物理距离，以替代传统可插拔方案，实现能效跃升。

■ **Rubin Ultra 机柜有望推动 M9 材料与 NPO 光引擎迎确定性增长。**Rubin Ultra 机柜以 144 颗 GPU（单颗 10.8TB/s 双向带宽）构建出了 1.5PB/s PB 级 Scale-up 网络，并采用 4-Canister 双层架构进行 Scale up 组网。第一层通过正交背板实现 Canister 内部无阻塞交换，考虑到 224G SerDes 对信号完整性的严苛要求，正交背板必须采用 M9 级超低损耗 CCL 材料；第二层采用 3:1 收敛设计组网，通过 72 颗 NVSwitch 与 648 颗 3.2T NPO 光引擎完成跨 Canister 光互连，GPU 与光引擎配比高达 1:4.5。

■ **CPO 交换机规模放量在即，核心供应商迎增长机遇。**英伟达 CPO 交换机产品矩阵抢先卡位，构筑 AI 网络新基建核心壁垒。Quantum X3450 作为全球首款量产 CPO 交换机，以 115.2T 带宽与可拆卸光引擎设计树立技术标杆；Spectrum X 平台（6810/6800）更以 102.4T-409.6T 梯度化带宽覆盖以太网生态，2026 年量产后将完善从 InfiniBand 到以太网的全场景布局。CPO 交换机规模化放量在即，光引擎、外部激光源、光纤连接单元等核心零部件需求将迎来爆发式增长，具备供应链卡位优势的国内龙头厂商有望率先受益，建议重点关注相关供应商投资机遇。

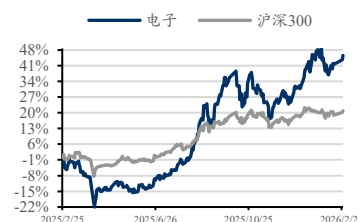
建议关注

M9 PCB 产业链：菲利华、东材科技、生益科技、胜宏科技、沪电股份、深南电路、东山精密等

CPO 产业链：致尚科技、长光华芯、源杰科技、仕佳光子、太辰光、炬光科技、罗博特科等

■ **风险提示：**算力互联需求不及预期，客户拓展及份额提升不及预期，产品研发及量产落地不及预期，行业竞争加剧

行业走势



相关研究

《2026 年端侧 AI 产业深度：应用迭代驱动终端重构，见证端侧 SoC 芯片的价值重估与位阶提升》

2026-02-23

《格局落定，价值归真：从周期波动走向技术溢价》

2026-02-06

内容目录

1. SerDes 代际跃迁，驱动算力互联介质升级	4
1.1. Serdes 持续升级，推动 GPU 互联带宽迭代增长	4
1.2. Serdes 速率提升，推动 CCL 升级	6
1.3. Serdes 功耗提升，推动光互联向近封装、共封装升级	7
2. Rubin ultra Scale up，正交背板+NPO 的双层网络结构	10
2.1. Kyber 架构机柜详解：Rubin ultra 互联带宽推演	10
2.2. Kyber 架构机柜详解：Canister 内部正交背板交换网络	11
2.3. Kyber 架构机柜详解：Canister 间 NPO 交换网络	12
3. 英伟达 CPO 交换机产品矩阵蓄势待发，供应链机遇全景透视.....	13
3.1. Quantum 5：InfiniBand 旗舰交换机的技术规格与集群部署	13
3.2. Spectrum 6：面向 AI 数据中心的以太网交换平台	15
3.3. CPO 供应链拆解：光引擎、硅光芯片及先进封装环节供应商梳理	17
4. 投资建议	19
5. 风险提示	20

图表目录

图 1: 英伟达 NVLink 代际演进	4
图 2: NVLink 3.0 带宽计算示意	5
图 3: Blackwell NVLink 5 架构示意	5
图 4: GB200 NVL 72 网络端口拆分	5
图 5: Vera Rubin NVL 144 (机柜总带宽 260TB/s)	6
图 6: Rubin ultra NVL 576 (机柜总带宽 1.5PB/s)	6
图 7: 各厂商 CCL 牌号比较及适用传输速率	6
图 8: 基于 NPO 光互联系统示意图	8
图 9: CPO 示意图	9
图 10: Ayar labs OIO 解决方案	10
图 11: 英伟达 AI 服务器迭代路径	11
图 12: Rubin Ultra NVL 576 架构	12
图 13: 英伟达 Quantum-X CPO 交换机	13
图 14: 英伟达 Quantum-X CPO 交换机 交换芯片	14
图 15: 英伟达 Quantum-X CPO 交换机 光引擎	15
图 16: 英伟达 Spectrum-X CPO 交换机 (左二)	15
图 17: 英伟达 Spectrum-X CPO 交换机 交换芯片 (左一)	16

1. SerDes 代际跃迁，驱动算力互联介质升级

1.1. Serdes 持续升级，推动 GPU 互联带宽迭代增长

算力芯片的互联带宽已成为衡量其系统级性能的核心指标，而决定这一指标上限的底层技术在于 SerDes 的代际演进。作为高速 IO 端口的关键组件，SerDes 负责将芯片内部并行数据流转换为高速串行信号，其单 lane 速率直接定义了 GPU 对外互联的带宽天花板。回顾英伟达 GPU 架构迭代路径，SerDes 速率呈现清晰的倍增规律：Ampere 架构采用 56Gbps SerDes 支撑 NVLink 3，Hopper 架构升级至 112Gbps SerDes 对应 NVLink 4，当前 Blackwell 架构则跃迁至 224Gbps SerDes 实现 NVLink 5 的带宽突破。这一技术演进并非简单的参数提升，而是直接决定了单机柜算力密度的理论上限——当 Rubin Ultra 机柜带宽达到 1.5PB/s 量级时，其底层必然对应着 SerDes 向 448G 乃至更高的跨越式升级。

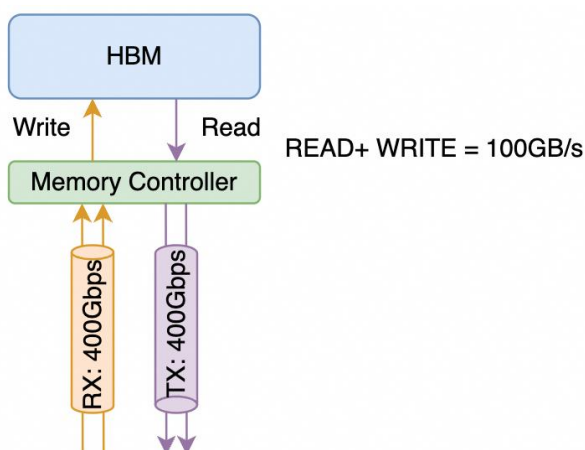
图1：英伟达 NVLink 代际演进

NVLink代际	NVLink 1.0	NVLink 2.0	NVLink 3.0	NVLink 4.0	NVLink 5.0
年份	2016	2017	2020	2022	2024
NVLink数量	4	6	12	18	18
通道数	32	48	48	36	36
单通道带宽	5 GB/s	6.25 GB/s	12.5 GB/s	25 GB/s	50 GB/s
调制方式	NRZ	NRZ	NRZ	PAM4	PAM4
总双向带宽	160 GB/s	300 GB/s	600 GB/s	900 GB/s	1800 GB/s

数据来源：鲜枣课堂，东吴证券研究所

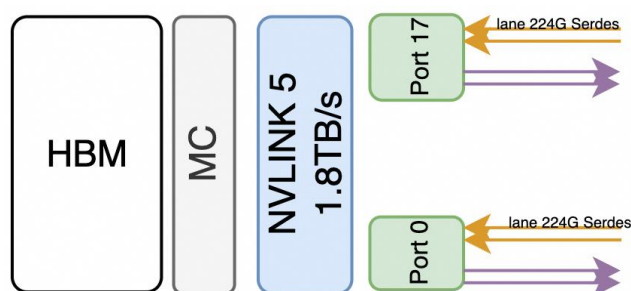
在分析 NVLink 带宽时，需首先厘清英伟达在计算口径与硬件定义上的技术混淆。一方面，GPU 计算侧通常沿用内存带宽的计量习惯，以字节每秒（Byte/s）为单位表述总线带宽；另一方面，NVLink Switch 及 IB/Ethernet 交换设备则采用网络设备视角，以比特每秒（bit/s）计量物理层传输速率。更深层的混淆在于 NVLink 的物理层定义：自 NVLink 3.0 起，英伟达采用"sub-link"作为基本物理单元，每个 sub-link 由 4 对差分信号线构成，同时包含独立的发送（TX）与接收（RX）通路。这与传统网络设备中"一个 400Gbps 接口指单方向 400Gbps 收发并发"的定义存在本质差异，导致同一物理接口在不同语境下呈现不同的带宽数值。

图2: NVLink 3.0 带宽计算示意



数据来源: Zartbot, 东吴证券研究所

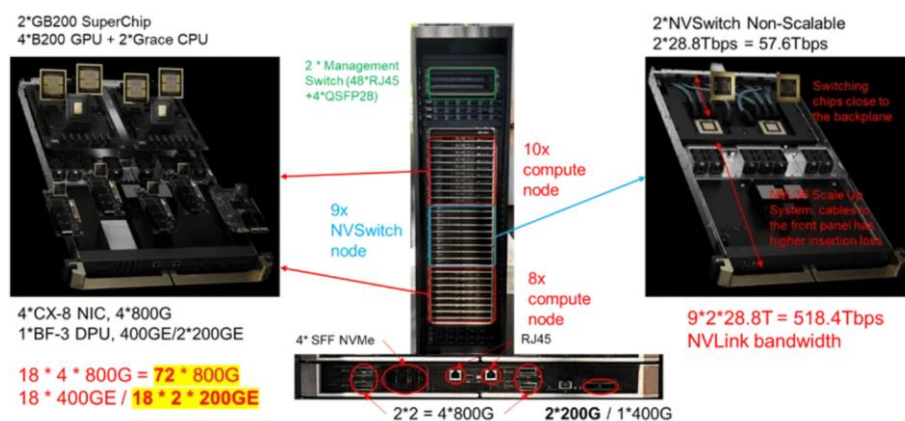
图3: Blackwell NVLink 5 架构示意



数据来源: Zartbot, 东吴证券研究所

以 Hopper 架构的 H100 为例可具体说明这一计算逻辑。H100 采用 112G SerDes, 经编码开销调整后单对差分线实际承载 100Gbps 有效数据。由于每个 sub-link 包含 4 对差分线且支持双向并发传输, 其单向网络带宽为 $4 \times 100\text{Gbps} = 400\text{Gbps}$, 而双向总带宽换算为字节单位则为 $400\text{Gbps} \times 2 \div 8 = 100\text{GB/s}$ 。但英伟达在 GPU 侧通常将收发合并表述为 50GB/s per sub-link (单向字节数), H100 共集成 18 个 sub-link, 故总互联带宽表述为 $50\text{GB/s} \times 18 = 900\text{GB/s}$ 。进入 Blackwell 时代, B200 采用 224G SerDes, 单对差分线速率提升至 200Gbps, 单个 sub-link 单向网络带宽达 800Gbps (即 $4 \times 200\text{Gbps}$), 对应 100GB/s 的双向字节带宽。18 个 sub-link 共同构建出 1.8TB/s ($100\text{GB/s} \times 18$) 的 NVLink 5 总带宽, 从网络设备视角等效于 9 个 400Gbps 单向接口的聚合能力。

图4: GB200 NVL 72 网络端口拆分

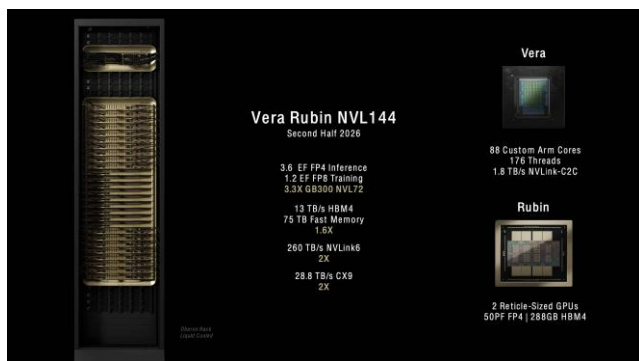


数据来源: 陆玉春《Nvidia AI 芯片演进解读与推演 (二)》, 东吴证券研究所

展望未来, SerDes 的迭代速度将进一步加快。根据英伟达官方技术路线图, Rubin 架构机柜级互联带宽较 GB200 提升 2 倍至 260TB/s, 而 Rubin Ultra 更将跃升 12 倍至 1.5PB/s。这一带宽跨越无法通过简单的 lane 数量堆叠实现, 必须依赖 SerDes 单 lane 速率的代际突破。基于当前 224G SerDes 支撑 1.8TB/s 单芯片带宽的技术基线推算, 要

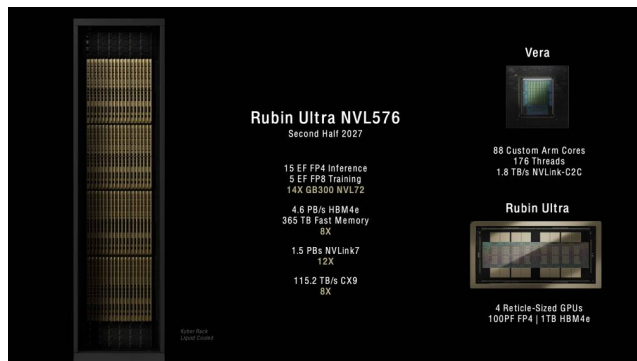
实现 PB 级机柜互联，SerDes 速率必然向 448G PAM4 乃至 896G PAM6 演进。这意味着下一代算力芯片的物理层设计将面临更严峻的信号完整性与功耗挑战，同时也为高速传输介质、光电封装技术产业链带来确定性的升级需求。

图5: Vera Rubin NVL 144 (机柜总带宽 260TB/s)



数据来源: 英伟达, 东吴证券研究所

图6: Rubin ultra NVL 576 (机柜总带宽 1.5PB/s)

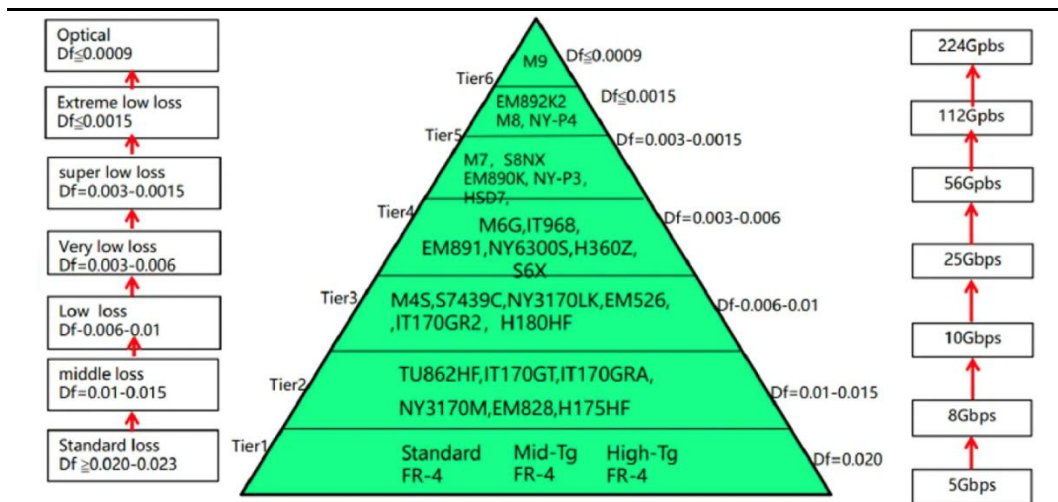


数据来源: 英伟达, 东吴证券研究所

1.2. Serdes 速率提升, 推动 CCL 升级

当 SerDes 速率沿着 56G-112G-224G 的路径持续攀升, 并朝着 Rubin 时代 448G 乃至 896G 演进时, 带宽增长的物理代价开始显现。如前文所述, Rubin Ultra 机柜要实现 1.5PB/s 的互联带宽, 意味着单链路速率必须突破当前 224Gbps 的物理极限。然而, 电信号传输遵循基本的物理规律: 224Gbps PAM4 调制信号的奈奎斯特频率已达 56GHz, 若进一步升级至 448G PAM4, 频率将飙升至 112GHz。在此频段下, 传统 M7/M8 级别覆铜板的介质损耗呈指数级增长, 信号在传输数英寸后便会衰减至无法恢复的程度。这意味着, 若不解决传输介质的物理瓶颈, 前述的带宽代际升级将无从落地, PCB 材料体系被迫迎来从 M7/M8 向 M9 (Df<0.001) 的强制性跃迁。

图7: 各厂商 CCL 牌号比较及适用传输速率



数据来源: SemiVision Research, 东吴证券研究所

而 CCL 材料升级的首要环节在于增强材料的革新。传统 PCB 采用 E-glass 玻纤布

($Dk \approx 6.6$, $Df \approx 0.001$)，其介电常数 (Dk) 和损耗因数 (Df) 在高频下表现不佳。为匹配 224G 以上 SerDes 需求，产业界正加速导入 Low Dk 玻纤布乃至熔融石英布。石英布凭借极低的介电损耗成为 M9 材料的核心选项，但其硬度高、编织难、与树脂结合力弱等工艺难点，导致目前仅少数头部覆铜板厂商具备量产能力。这种材料壁垒直接决定了 M9 覆铜板的供应稀缺性。

此外，树脂基体体系的变革则是 M9 材料实现低损耗的根本保障。传统环氧树脂 (Epoxy) 因含羟基、环氧基等极性基团， Df 值较高，难以满足 112G 以上 SerDes 的低损耗需求。在 AI 算力及高频通信领域，其正逐步被聚苯醚 (PPO)、碳氢树脂乃至苯并环丁烯 (BCB) 等低极性材料替代。特别是 BCB 树脂，其 Df 值可低至 0.0008，且玻璃化转变温度 $>350^\circ\text{C}$ ，成为支撑 224G SerDes 及 CPO 封装的前沿候选材料。碳氢树脂体系则通过引入氢化环烯烃共聚物与聚丁二烯共混改性，配合球形硅微粉的界面极化抑制技术，在 M8/M9 级别实现 $Df < 0.001$ 的平衡性能。这些特种树脂的合成涉及分子量精准控制、低极性交联剂选择及填料表面改性，构成了覆铜板厂商的核心 Know-how。

最后，CCL 铜箔表面处理技术的同步升级同样关键。高频信号遵循趋肤效应，在 224Gbps 速率下，信号在铜箔表面的趋肤深度仅约 $0.4\mu\text{m}$ ，铜箔表面粗糙度 (R_z) 若过大将显著增加导体损耗。传统 HTE (高延伸率) 铜箔 R_z 值约 $3\text{-}5\mu\text{m}$ ，已无法满足需求；产业正向 HVLP4 乃至 HVLP5 铜箔迁移。更先进的 VLP (Very Low Profile) 与 ULP (Ultra Low Profile) 铜箔通过特殊晶粒结构控制与表面粗化处理工艺，在保障与树脂结合强度的同时，将表面粗糙度降至亚微米级，确保 56GHz 以上信号的传输完整性。

从产业投资视角看，SerDes 向 448G 演进将引发 PCB 材料体系的代际替换潮。我们预计 M9 级别 CCL 在 AI 服务器 PCB 中的渗透率将从 2026 年快速提升，驱动覆铜板行业价值量重构。建议重点关注具备石英纤维、石英布编织技术、HVLP4/5 铜箔量产能力，以及 BCB/碳氢树脂配方的上游材料龙头。这一技术升级并非简单的工艺改进，而是由算力芯片物理层速率瓶颈决定的必然选择，产业链准备度与认证壁垒将决定未来两年的竞争格局。

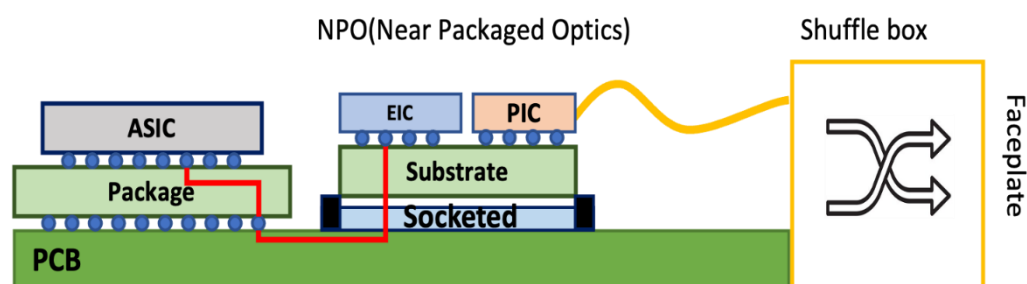
1.3. Serdes 功耗提升，推动光互联向近封装、共封装升级

SerDes 速率的指数级提升在突破带宽瓶颈的同时，也带来了严峻的功耗挑战。缩短光电转换点与交换芯片之间的电气距离，减少甚至消除高功耗 DSP 的使用，成为光互联技术演进的核心逻辑。当前 800G 光模块采用 $8 \times 100\text{G}$ SerDes 架构，单模块功耗已达 12-18W；随着 SerDes 向 200G/lane 演进，1.6T 光模块 ($8 \times 200\text{G}$) 功耗已飙升至 25-30W，其中 DSP (数字信号处理器) 用于补偿信道损耗的功耗占比超过 50%。更严峻的是，当 SerDes 速率向 448G 乃至更高迭代时，电信号在 PCB 走线及连接器中的高频损耗呈指数级增长——根据 OIF (光互联论坛) 数据，448G PAM4 信号在标准 PCB 上的插入损耗可达 20-50dB，必须依赖更高阶的 DSP 算法进行补偿，这将导致单通道功耗突破 15W，3.2T 光模块整体功耗将接近 40W。业界预计届时 SerDes 在交换芯片中的功耗

占比将超过 40%，热流密度高达 $50\text{W}/\text{cm}^2$ ，传统风冷散热已触及物理极限。因此，单纯依靠工艺微缩和算法优化已无法破解功耗困局，缩短光电转换点与交换芯片之间的电气距离，减少甚至消除高功耗 DSP 的使用，成为光互联技术演进的核心逻辑。

NPO (Near Packaged Optics, 近封装光学) 作为过渡性方案，率先在产业化路径上取得突破。该技术将光引擎通过 LGA 连接器直接部署在交换机板上，与交换芯片的物理距离缩短至 150mm(符合 OIF 标准)，远小于传统可插拔模块 15-30cm 的走线长度。电气路径的缩短显著降低了高频信号衰减，使得光引擎完全省去高功耗 DSP 芯片，采用线性直驱架构，仅保留 Driver 和 TIA 等模拟器件，从而相较传统可插拔模块降低 50% 以上功耗，更为重要的是，NPO 保留了光引擎的可拆卸性，支持热插拔维护，避免了与交换芯片绑定封装带来的良率风险及“故障需更换整机”的运维难题。因此，NPO 成为云服务商现阶段优先落地的技术方案——如阿里云在 UPN512 超节点（512 颗 xPU 全互联架构）中明确采用 NPO 作为核心使能技术，并已成功点亮全球首款 3.2T NPO 光模块。

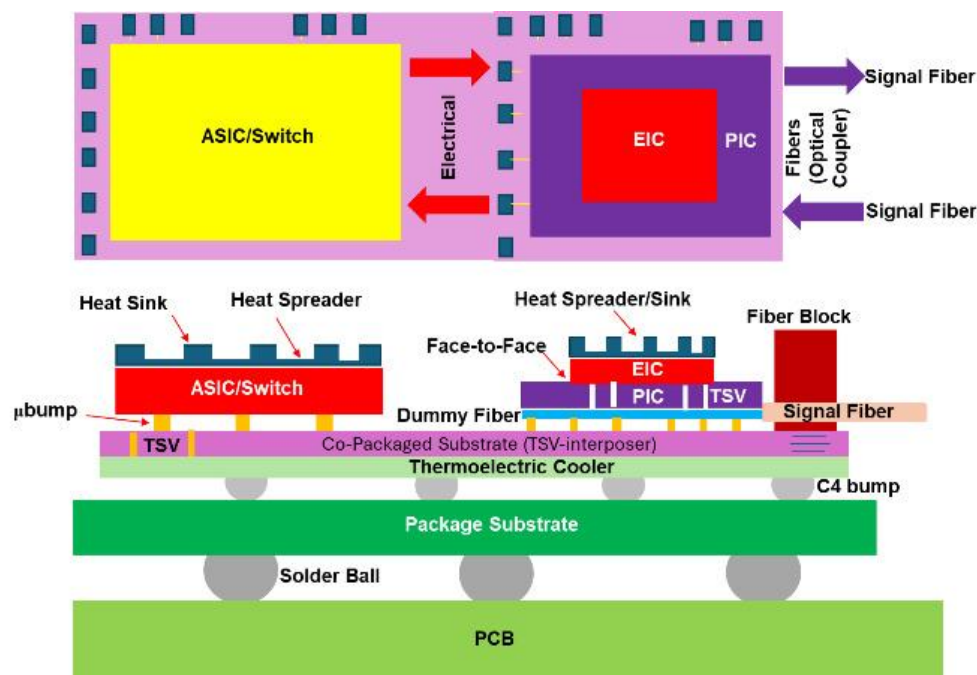
图8: 基于 NPO 光互联系统示意图



数据来源：讯石光通信网，东吴证券研究所

CPO 则代表了中期内的终极解决方案。该技术将光引擎与交换芯片共同封装在同一 IC 载板或硅中介层上，电气连接距离进一步压缩至 50mm 以内（符合 OIF 标准），部分基于 FOWLP 的先进方案甚至实现亚毫米级的极短互连。根据博通与 Meta 的联合实测数据，CPO 方案可将 800G 光互联功耗从传统可插拔模块的 15W 降至 5.4W，系统整体能耗降低 65% 以上。英伟达技术数据显示，CPO 可将每端口功耗从 30W 降至 9W，信号完整性提升 64 倍。然而，CPO 面临光引擎与交换芯片绑定封装导致的良率耦合问题——一旦光引擎失效需整体更换 ASIC 模组，维护成本高昂，因此业界预计将在 2026-2027 年随着封装工艺成熟逐步商用。

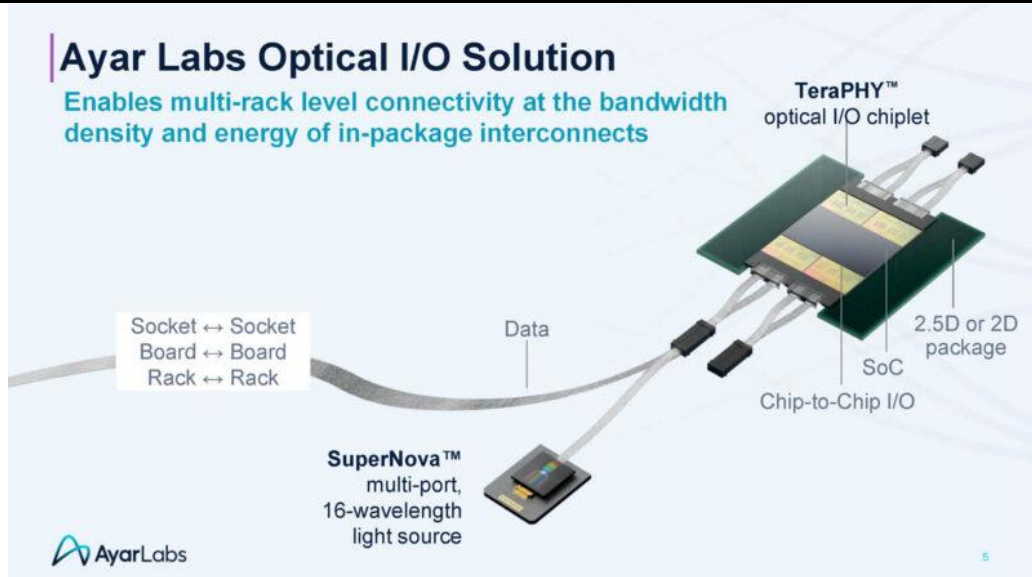
图9: CPO 示意图



数据来源: Journal of Microelectronics & Elect Pkg, 东吴证券研究所

OIO (Optical I/O, 光学 I/O) 则是面向 Rubin 及更远期架构的终极愿景。该技术将光收发功能直接集成至计算芯片的封装基板或硅中介层上, 与 HBM 及计算裸片通过 TSOV 实现 3D 堆叠, 彻底取代传统的电 I/O 接口。OIO 可将光引擎、GPU、HBM 置于同一封装内, 实现芯片级光互连。根据中国信通院数据, 相比传统可插拔方案, OIO 可将数据传输带宽提升 7 倍, 功耗降低至 1/5, 尺寸缩小至 1/12。Intel 的 OCI (Optical Compute Interconnect) Chiplet 已验证该架构可行性, 其 4Tbps 双向传输速率与 <5pJ/bit 的能效显著优于电 I/O 方案。尽管目前 OIO 仍处于实验室向产业化过渡阶段, Yole 预测其商业生态需 5 年以上方能完全爆发 (Yole 预计 2033 年市场规模达 23 亿美元), 但已代表了"光进铜退"在芯片尺度的终极形态。

图10: Ayar labs OIO 解决方案



数据来源: Ayarlabs, 东吴证券研究所

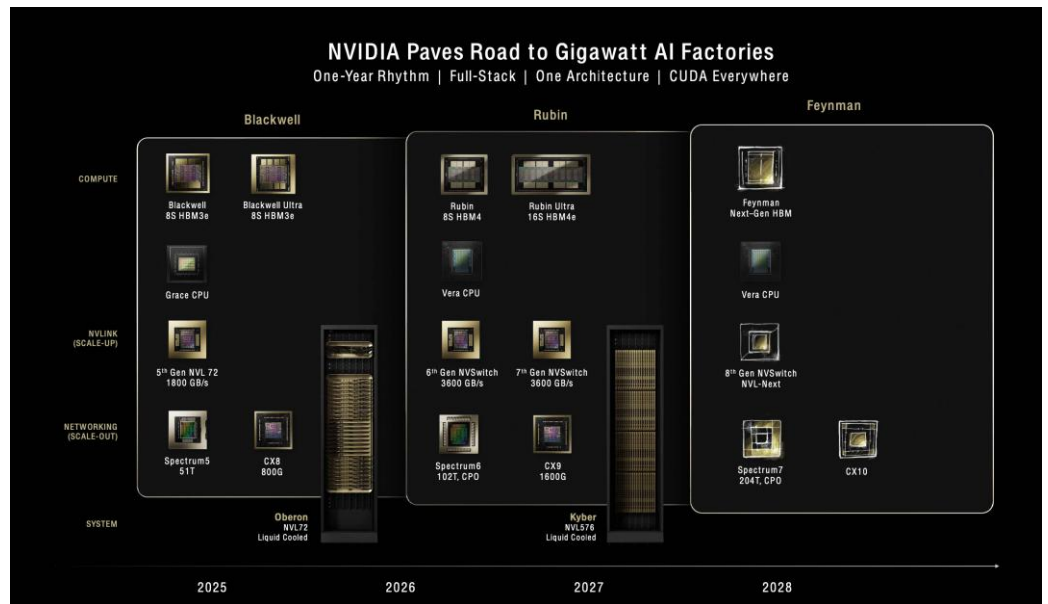
从 NPO 到 CPO 再到 OIO, 光互联技术正沿着"板级近封装→载板共封装→芯片内集成"的路径持续演进, 以应对 SerDes 向 448G/896G 升级带来的功耗与带宽密度双重挑战。

2. Rubin ultra Scale up, 正交背板+NPO 的双层网络结构

2.1. Kyber 架构机柜详解: Rubin ultra 互联带宽推演

NVIDIA Rubin Ultra 作为下一代 AI 数据中心 GPU 旗舰产品, 在互联带宽方面实现了跨越式升级, 标志着 AI 基础设施正式进入 PB 级全互联时代。根据 NVIDIA 在 GTC 2025 大会公布的技术规格, Rubin Ultra 机柜总带宽较上一代 GB200 增长 12 倍, 达到 1.5 PB/s (1555.2 TB/s)。这一突破性的带宽能力由 144 颗 GPU 共同提供, 每颗 GPU 封装集成 4 个 dies, 总计 576 个计算 die, 对应单颗芯片双向互联带宽达到 10.8 TB/s。如此高密度的算力与带宽配置, 对 Scale-up 网络架构提出了前所未有的工程挑战, 或将采用双层网络拓扑以平衡极致性能、成本控制与物理可实现性。

图11：英伟达 AI 服务器迭代路径



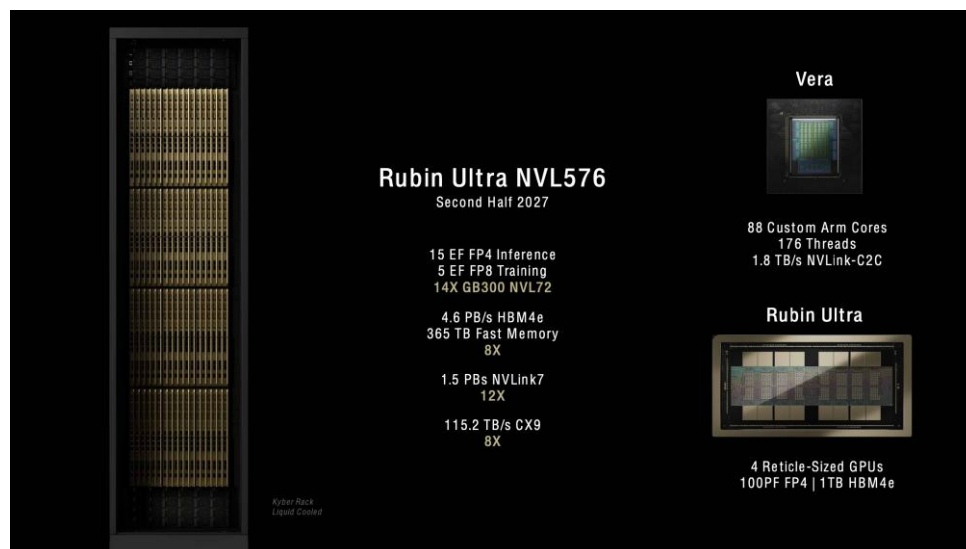
数据来源：英伟达，东吴证券研究所

从交换侧来看，Rubin Ultra 将配备第七代 NVSwitch 芯片。根据官方 PPT 披露的技术细节，单颗第七代 NVSwitch 芯片的交换容量达到 3600 GB/s (3.6 TB/s)。考虑到 Rubin Ultra 机柜采用 4 个 Canister 堆叠构成的物理形态，这种模块化设计不仅优化了机柜的空间布局与散热管理，也为分层网络架构提供了天然的物理边界。假设第七代 NVSwitch 的 SerDes 仍沿用成熟的 224G 技术路径，radix 端口数维持 72 配置，整个 Scale-up 网络必须进行双层组网架构设计：第一层负责 Canister 内部的高速电交换，第二层负责跨 Canister 的光互连。这种分层设计不仅符合信号完整性的工程约束，也通过电光混合方案在带宽密度与传输距离之间取得了最优平衡。

2.2. Kyber 架构机柜详解：Canister 内部正交背板交换网络

第一层网络负责 Canister 内部计算托盘与交换托盘之间的高速互联，是实现机柜级全互联的基础单元。按照历次 GTC 大会展示的技术方案及产业链验证，这一层的信号传输介质采用正交背板 PCB。正交背板架构通过将计算板与交换板垂直插接，最大限度地缩短了高速信号走线长度，有效降低了传输损耗和信号延迟，为高带宽、低延迟的片间通信提供了理想的电气环境。考虑到 224G SerDes 的高频信号完整性要求，正交背板所使用的 CCL 材料需达到 M9 级别。M9 级超低损耗材料具有优异的介电常数(Dk)稳定性和极低的介电损耗(Df)，能够有效降低高频信号在 PCB 走线中的插入损耗，确保 224G PAM4 信号在正交背板上的传输质量，从而支撑 Canister 内部全带宽无阻塞交换。

图12: Rubin Ultra NVL 576 架构



数据来源：英伟达，东吴证券研究所

从带宽层面进行定量分析，单颗 Rubin Ultra 芯片具备 10.8 TB/s 的双向互联带宽，单个 Canister 包含 36 颗 Rubin Ultra 芯片。考虑到交换芯片的容量规划通常以单向带宽为基准，整个 Canister 的总单向带宽需求达到 194.4 TB/s。基于第七代 NVSwitch 单颗 3.6 TB/s 的交换容量，第一层网络需配置 54 颗第七代 NVSwitch 芯片方可满足该 Canister 内部的无阻塞交换需求。这 54 颗交换芯片通过正交背板与 36 颗计算芯片实现高密度的端口映射，构建出 Canister 内部的全互联拓扑。这种高密度的交换芯片配置不仅提供了充足的交换容量，还通过缩短信号路径降低了通信延迟，为 Canister 内部 36 颗 GPU 之间的高速通信提供了坚实的电气基础，确保训练过程中梯度同步等关键操作能够以最高效率完成。整个机柜共包含 4 个 Canister，因此第一层网络总共需要 4 块正交背板，216 颗第七代 NVSwitch 完成内部互联。

2.3. Kyber 架构机柜详解：Canister 间 NPO 交换网络

第二层网络用于实现 4 个 Canister 之间的跨节点互联，这是构建完整 144 GPU 机柜级全互联的关键环节，也是实现 1.5 PB/s 机柜总带宽的必经之路。由于 4 个 Canister 在物理上呈堆叠或并排布局，跨 Canister 的信号传输距离超出了正交背板铜互连的有效范围，因此需采用光互连技术。根据 Global Technology Research，Canister 之间的网络带宽采用 3:1 的收敛比设计，这一工程权衡在极致性能与成本控制之间取得了合理平衡，既保证了机柜内任意 GPU 间的可达性，又通过适度的带宽收敛控制了光引擎的数量和系统总成本。

要支撑整个机柜双向 1.5 PB/s (1555.2 TB/s) 的带宽互联需求，第二层网络所需的 NVSwitch 芯片数量可通过严谨的数学计算得出：首先将双向总带宽 1555.2 TB/s 除以 2 转换为单向带宽 777.6 TB/s，再除以 3 考虑 3:1 的收敛比得到 259.2 TB/s，最后除

以单颗芯片 3.6 TB/s 的交换容量，即 $1555.2 / 2 / 3 / 3.6 = 72$ 颗第七代 NVSwitch 芯片。第二层网络中这 72 颗 NVSwitch 的总交换带宽达到 259.2 TB/s，换算为比特率即 2073.6 Tbps。这些交换芯片充当着 Canister 间通信的汇聚与转发节点，将来自第一层网络的电信号转换为适合长距离传输的格式。

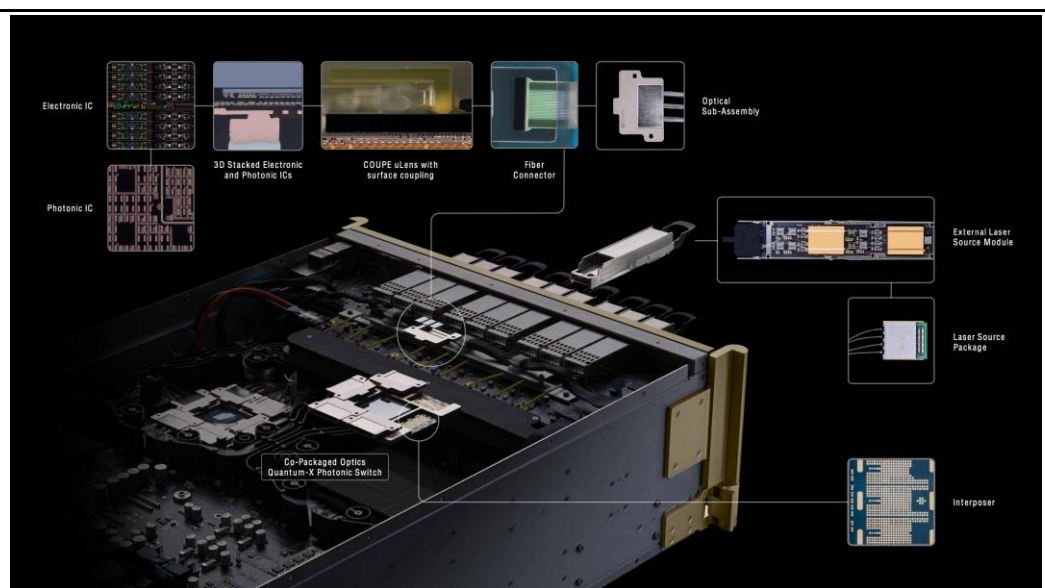
根据 Global Technology Research 的技术披露，第二层网络中 Canister 之间的互联介质采用 3.2 Tbps 速率的光引擎（NPO，Near-Package Optics）。基于上述总带宽需求 2073.6 Tbps，需配置 648 颗 3.2 T 光引擎以完成第二层网络的互联。这种近封装光学方案将光模块紧密封集成在交换芯片附近，极大地缩短了电信号走线，降低了功耗和信号完整性风险。光互连方案通过将电信号转换为光信号，突破了铜缆在传输距离和带宽密度上的物理限制，使得 4 个 Canister 能够在机柜范围内实现高速、低延迟的协同计算，构建出真正意义上的机柜级统一加速计算资源池。这一配置对应 GPU 与光引擎的比例为 144:648，即 1:4.5。该配置方案不仅满足了 1.5 PB/s 机柜带宽的技术指标，也为未来超大规模 AI 模型训练与推理工作负载提供了可扩展、高可靠、高能效的互联基础设施，代表了当前 AI 数据中心 Scale-up 网络架构的巅峰工程水平。

3. 英伟达 CPO 交换机产品矩阵蓄势待发，供应链机遇全景透视

3.1. Quantum 5: InfiniBand 旗舰交换机的技术规格与集群部署

英伟达 Quantum X800-Q3450 作为全球首款 CPO 交换机已于 2025 年下半年正式面市，其技术架构代表了数据中心网络演进的重要里程碑。该设备整机配备 144 个物理 MPO 端口，可灵活配置为 144 个 800G 逻辑端口或 72 个 1.6T 逻辑端口，总交换带宽高达 115.2T，充分展现了 CPO 技术在端口密度上的突破性进展。

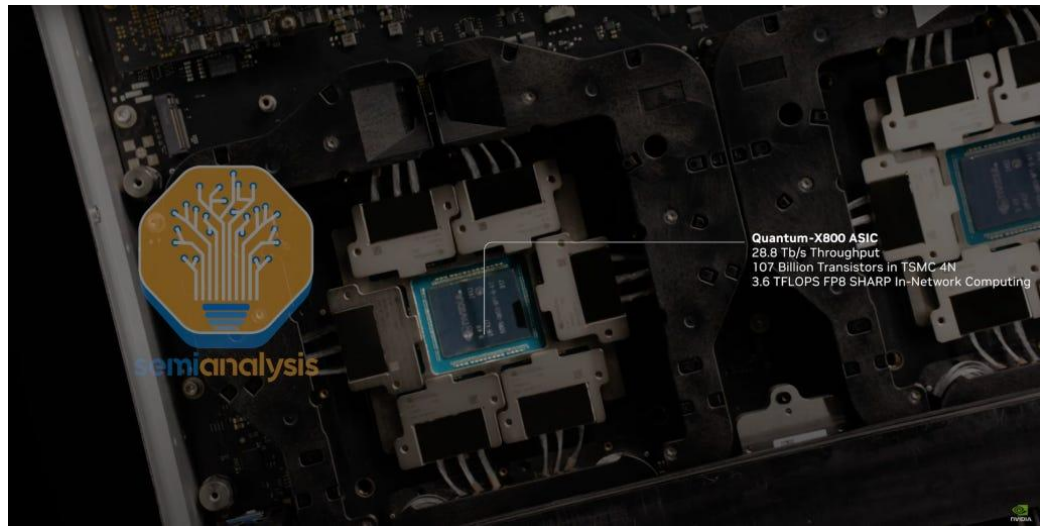
图13：英伟达 Quantum-X CPO 交换机



数据来源：英伟达，东吴证券研究所

从核心交换架构来看，Quantum X800-Q3450 采用四颗 Quantum-X800 ASIC 芯片成多平面交换配置，每颗芯片提供 28.8 Tbit/s 的交换能力。这一多平面设计的精妙之处在于，每个物理端口均与四颗交换芯片同时建立连接，使得任意端口的数据流均可通过四颗独立交换芯片分散至四条 200G 通道进行并行处理，从而在实现超高聚合带宽的同时，显著提升了系统的冗余能力与负载均衡性能。这种架构设计有效解决了传统单芯片交换方案在高基数场景下的瓶颈问题，为大规模 AI 集群的横向扩展提供了坚实的网络基础。

图14：英伟达 Quantum-X CPO 交换机 交换芯片



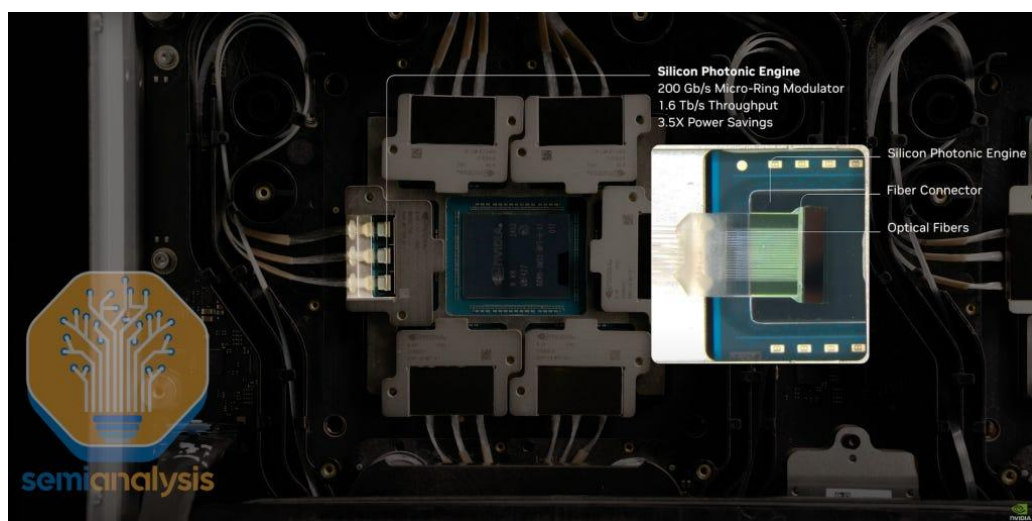
数据来源：Semi analysis，英伟达，东吴证券研究所

在光电集成层面，每颗 Quantum-X800 ASIC 被六个可拆卸的光学子组件（OSA，Optical Sub-Assembly）环绕，每个子组件内置三个光学引擎，单颗 ASIC 共计配置 18 个光学引擎。每个光学引擎提供 1.6 Tbit/s 的光学带宽，因此每颗 ASIC 的光学总带宽达到 28.8 Tbit/s，与电交换能力完美匹配。值得特别关注的是，这些光学子组件采用可拆卸设计，从严格意义上讲这使得该设备更接近 NPO（Near-Packaged Optics，近封装光学）的范畴，而非纯粹的 CPO 架构。尽管如此，可拆卸设计带来的额外信号损耗对整体性能的影响微乎其微，却大幅提升了设备的可维护性与升级灵活性，这一权衡体现了英伟达在工程实现与商业可行性之间的深思熟虑。

深入至光学引擎的内部架构，每个 1.6T 光引擎集成 8 条电气与光学通道，形成了完整的光电信号转换链路。在电气侧，采用 200G PAM4 SerDes 技术驱动高速电信号；在光学侧，则通过 8 个微环调制器（MRM，Micro-Ring Modulator）运用 PAM4 调制技术，实现每个调制器 200G 的传输速率。每个光学引擎内部集成两颗异构芯片：光子集成电路（PIC）基于成熟的 N65 工艺节点构建，集成调制器、波导和探测器等光学组件，这些器件无法从工艺微缩中获得性能增益，通常在较大的几何尺寸下表现更优；电子集成电路（EIC）则采用先进的 N6 工艺节点制造，包含驱动器、跨阻放大器和控制逻辑等电路，这些组件能够显著受益于先进制程带来的更高晶体管密度与更优能效表现。两颗

芯片通过台积电的 COUPE 平台进行混合键合，实现了光子域与电子域之间超短距离、高带宽的互连，有效降低了信号完整性损耗与系统功耗。

图15: 英伟达 Quantum-X CPO 交换机 光引擎

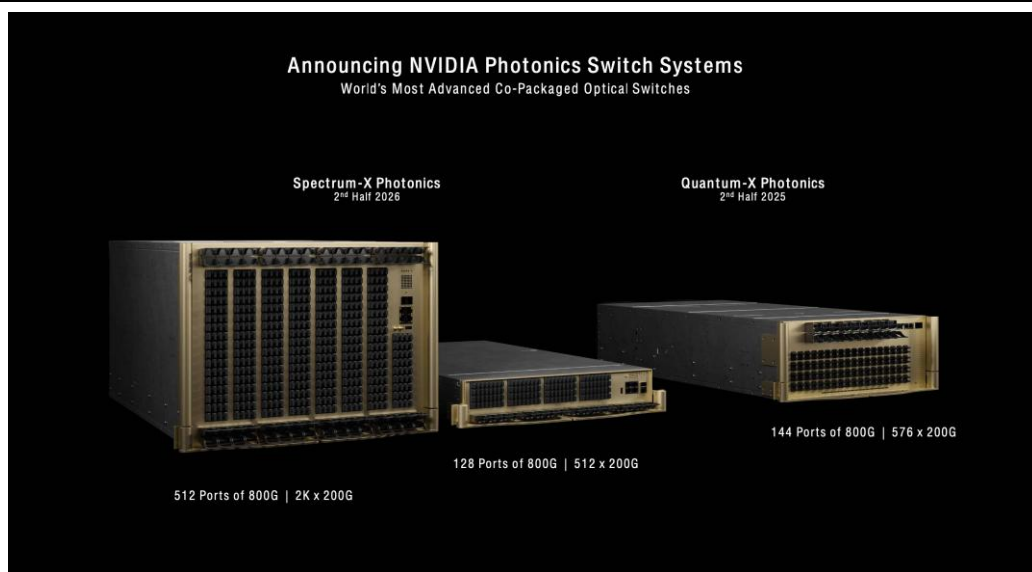


数据来源: Semi analysis, 英伟达, 东吴证券研究所

3.2. Spectrum 6: 面向 AI 数据中心的以太网交换平台

继 Quantum X800-Q3450 之后，英伟达计划于 2026 年下半年推出 Spectrum-X Photonics 系列 CPO 交换机产品,该系列将包含两款独立配置: Spectrum 6810 提供 102.4T 聚合带宽，以及更大规模的 Spectrum 6800——后者通过采用四个独立的 Spectrum-6 多芯片模块（MCM, Multi-Chip Module）实现高达 409.6T 的聚合带宽。这一产品矩阵的推出，标志着英伟达在 CPO 领域形成了从 InfiniBand 到以太网的完整技术布局，进一步巩固其在数据中心网络市场的领先地位。

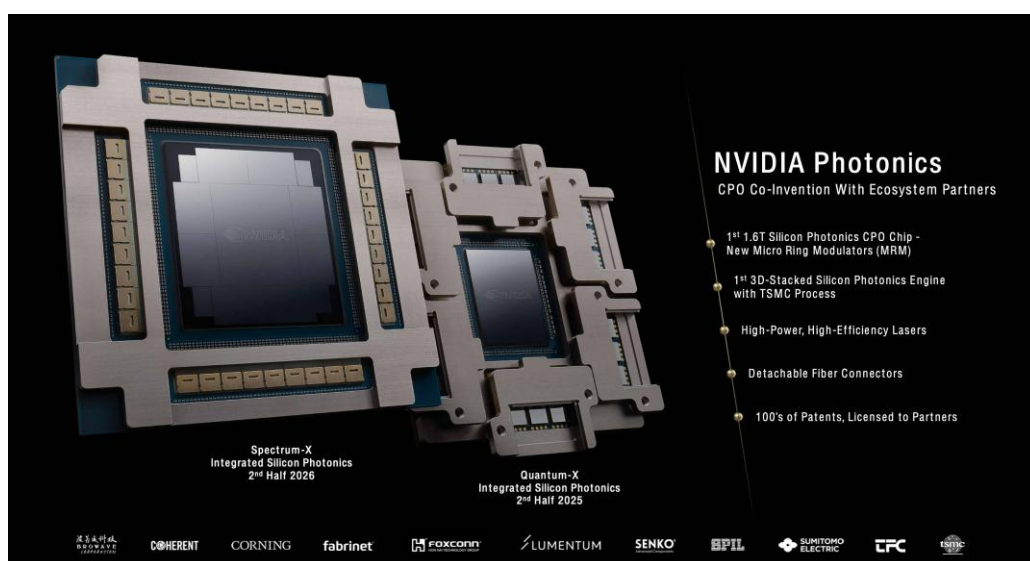
图16: 英伟达 Spectrum-X CPO 交换机（左二）



数据来源: Semi analysis, 英伟达, 东吴证券研究所

从芯片架构层面分析，Spectrum-X Photonics 与 Quantum X800-Q3450 存在本质差异。Quantum X800-Q3450 采用多平面配置，将四颗独立的单晶片交换机 ASIC 分别封装后连接至物理层，每颗 ASIC 内嵌 28.8T 交换能力及必要的 SerDes 与其他电子元件。与之相对，Spectrum-X 光子交换机采用多芯片模块（MCM）设计，其核心为具有更大光罩尺寸的 102.4T 交换机 ASIC，周围环绕八颗 224G SerDes I/O 小芯片，每侧各布置两颗。这种 MCM 架构的优势在于，通过将大型交换芯片与专用 I/O 芯片解耦，既突破了单芯片光罩尺寸的限制，又为 SerDes 预留了更多的布线边缘空间与布局面积，从而在同等封装尺寸下实现了更高的集成度与性能密度。

图17：英伟达 Spectrum-X CPO 交换机 交换芯片（左一）



数据来源：英伟达，东吴证券研究所

在光学引擎配置方面，每个 Spectrum-X 光子多芯片模块交换机封装将集成 36 个光引擎，这些光引擎采用英伟达第二代光学引擎技术，单引擎提供 3.2T 带宽，每个光引擎包含 16 条光通道，每条通道速率为 200G。值得注意的是，实际工作的光引擎为 32 个，额外配置的 4 个光引擎用于冗余备份，以应对光引擎焊接在基板上不易更换的可靠性挑战。这一设计体现了英伟达对 CPO 系统可维护性的深刻洞察——虽然 CPO 技术将光学器件与交换芯片紧密集成以缩短电互连距离，但也带来了传统可插拔光模块所不具备的维修难题，冗余设计成为平衡性能与可靠性的必要妥协。

从系统级架构来看，每个 I/O 芯片组提供 12.8T 总单向带宽，包含 64 条 SerDes 通道，并与 4 个光引擎分别对接。正是这种高度并行的设计，使得 Spectrum-X 能够提供远高于 Quantum-X Photonics 的总带宽能力。Spectrum-X 6810 交换机箱采用单个上述交换单元，可提供总计 102.4T 的交换容量，定位于中等规模的 AI 集群部署场景。而规模更大的 Spectrum-X 6800 交换机箱则是一款高密度旗舰产品，通过使用四块 Spectrum-X 交换单元，并以多平面配置连接至外部物理端口，实现了总计 409.6T 的惊人带宽。与 Quantum X800-Q3450 类似，Spectrum-X 6800 同样采用内部分支连接设计，将每个端口物理连接至全部四颗 ASIC，确保任意端口均可充分利用多平面架构的聚合带宽与冗余

能力。

Spectrum-X 系列的推出具有重要的战略意义。一方面，它将以太网生态引入 CPO 技术范畴，使得更广泛的云服务商与企业数据中心能够受益于光电共封装带来的功耗与密度优势；另一方面，102.4T 与 409.6T 两档带宽配置形成了清晰的产品梯度，可分别对标不同规模的 AI 训练集群与高性能计算场景。随着 AI 大模型参数规模的持续膨胀，数据中心内部的东西向流量呈现指数级增长，传统可插拔光模块在功耗、密度与成本方面的瓶颈日益凸显，Spectrum-X 系列的量产交付将为这一挑战提供系统性的技术解决方案。

3.3. CPO 供应链拆解：光引擎、硅光芯片及先进封装环节供应商梳理

CPO 交换机的零部件供应链涵盖光学、电子与结构件等多个细分领域，其技术门槛与价值分布呈现出明显的层次化特征。以英伟达 X800-Q3450 CPO 交换机为例，该设备总吞吐量达 115.2T，专为大规模 AI 集群的横向扩展网络设计。其初始版本搭载 72 个光引擎，每枚运行速度为 1.6Tbit/s；业界预计后续版本将升级为 36 个光引擎，每枚速率达 3.2Tbit/s。

在激光源领域，X800-Q3450 采用 18 个 ELS 模块作为激光源，每个模块包含 8 个连续波（CW）DFB 激光芯片，CPO 系统需要使用相对较高功率的激光源，每个 CW-DFB 芯片的功率输出约为 350mW。根据 Semi analysis，具备连续波激光单元量产能力的关键行业参与者包括博通、古河电工、Lumentum、Coherent、源杰科技、长光华芯与仕佳光子等。从竞争格局来看，Lumentum、Coherent、古河电工和博通的定价通常高于中国供应商，这主要源于其在高端光通信市场的长期技术积累与品牌溢价。Semi analysis 预计 Lumentum 将成为英伟达首批 CPO 交换机产品的独家供应商，Coherent 可能于 2026 年末作为第二供应商加入供应链体系。中国制造商未来可能迎来发展机遇，因为连续波激光源总体被视为相对标准化和商品化的部件，但构建 CPO 应用所需的高功率激光源仍存在一定的技术壁垒，尤其是在波长稳定性、功率效率与长期可靠性方面。

光纤连接单元（FAU，Fiber Array Unit）是实现光纤与光引擎精确耦合的关键被动元件，其对准精度直接决定系统的光学性能与插入损耗。X800-Q3450 上的每个 1.6T 光引擎配置一个包含 20 根光纤的连接单元，其中 8 根用于发射，8 根用于接收，4 根用于外部激光器。每套系统的光纤总数达 1440 根，其中包括 1152 根用于发射/接收的光纤。根据 Semi analysis，FAU 领域的领先企业包括天孚通信、Senko 与上詮。Semianalysis 表示天孚通信极有可能成为 X800-Q3450 CPO 交换机 FAU 的主要供应商，而 Senko 则被视为 NVIDIA Spectrum-X CPO 和博通 Tomahawk 6 CPO 系统的最有力候选供应商，上詮预计将更专注于英伟达大型 CPO 解决方案领域。天孚通信的核心竞争力在于其强大的制造能力与中国本土大量熟练且具成本效益的劳动力资源，此外天孚通信约在三年前就开始与英伟达在共封装光学设计方面展开合作，这一早期布局正逐步转化为供应链优势。Senko 则拥有标志性的 SEAT（Senko Elastic Average Technology）平台，可为 CPO 系统

提供可分离的 FAU 解决方案，公司正与 GFS 在边缘耦合技术方面紧密合作，将 Senko 反射镜直接集成到晶圆槽中，实现边缘耦合的晶圆级测试能力。

光纤交换箱 (Fiber Shuffle Box) 是 CPO 系统中负责光纤布线与端口映射的关键结构件。以 X800-Q3450 为例，其光学引擎引出上千根光纤，需要通过交换箱进行布线整理并导向目标端口。传统上这种光纤对位工序依赖人工操作，但部分行业领军企业已研发出能实现更高精度与效率的自动对位专用设备。交换箱的主要物料清单组件包括 MT 陶瓷插芯和光纤。根据 Semi analysis，太辰光是光纤配线箱行业的领军企业，公司已研发出用于光纤配线箱内光纤自动对准的自动化设备。太辰光的主要客户是康宁，两家公司通常协同服务客户，康宁负责帮助英伟达、博通等客户设计其 CPO 解决方案中的光纤网络，并将配线箱部分外包给天孚通信。

MPO 连接器位于交换箱的外围区域，负责将交换箱内部的光纤与外部端口相连，随后光纤 MPO 线缆可插入该端口，实现交换机与远端设备的光互连。MPO 连接器的制造过程主要涵盖注塑成型、真空灌胶以及单元内 MT 插芯的穿纤环节。对于 X800-Q3450 CPO 交换机而言，总共需要 144 个 MPO 连接器。根据 Semi analysis，能够生产 MPO 连接器的厂商包括 US Conec、太辰光、Senko、Broadex 与 Optec，这些厂商通常需要与光纤网络合约制造商合作，将其组件纳入整体网络设计方案中。

MT 插芯是 FAU、交换箱以及 MPO 连接器中至关重要的组件，其作用是以并行方式对齐多根光纤。根据 Semi analysis，目前全球多家企业具备 MT 插芯的生产能力，包括 US Conec、太辰光、Senko、福可喜玛、FOCI、住友电工与天孚通信。MT 插芯的制造工艺本身并不特别复杂，但要达到所需的精度和耐用性仍需要深厚的工程积累，各家企业的竞争焦点主要体现在注塑成型技术上，通过优化插芯制造工艺来最大限度地降低光纤连接时的插入损耗。US Conec 拥有超过三十年的技术研发经验，Semi analysis 预计公司将成为英伟达 Q3450 CPO 系统的主要供应商之一；福可喜玛具备强大的模具设计与制造能力，能够以具有竞争力的价格生产高质量 MT 插芯；FOCI、天孚通信与太辰光主要为其内部 FAU 和交换箱生产 MT 插芯，通过垂直整合生产流程提升质量控制与成本效益。

CPO 交换机的制造供应链涵盖光电芯片制造、先进封装与测试设备等多个高附加值环节，其技术复杂度与资本投入门槛显著高于传统光模块产品。在光电制造环节，台积电将在光子集成电路、电子集成电路及其集成，以及 CPO 系统的制造中发挥至关重要的作用，其 COUPE 平台有望成为下一代 CPO 端点的首选解决方案。Global Foundries 与 Tower 同样是具备硅光能力的代工厂商，但二者缺乏尖端 CMOS 技术与先进封装能力，这限制了其为未来更高带宽光电引擎提供制造服务的能力边界。台积电凭借其在先进制程、异构集成与封装技术方面的综合优势，确立了在 CPO 光电制造领域的核心地位。

在先进封装环节，外包半导体封测 (OSAT) 厂商将专注于后端工艺，包括光引擎

封装、光引擎测试以及系统级封装(激光器与耦合器的集成及测试)。根据 Semi analysis，日月光、安靠科技与讯芯科技是该领域的主要供应商。其中日月光作为英伟达供应链中的长期合作伙伴，未来将参与 Rubin 平台共封装光学系统的生产；讯芯科技则与博通保持着密切的合作关系。其他值得关注的厂商包括菲尼萨网络、天孚通信与鸿海精密。菲尼萨长期为英伟达内部光模块单元提供模块组装服务，目前正积极布局光引擎封装测试及全系统集成能力，该公司与迈艾世嘉、鸿海精密同被视为博通共封装光学系统组装的潜在合作伙伴。天孚通信在过去三到四年间与英伟达就共封装光学设计保持紧密合作，并将成为 FAU 的主要供应商，该公司已在中国苏州投资建设先进封装工厂，彰显其意在 CPO 供应链中占据更重要地位的雄心。

在电光测试设备领域，各封装测试服务商采用电光测试工具来确保系统可靠性，但当前行业仍处于发展初期，尚未就光子引擎的标准化测试方法达成共识，各厂商正在开发多种解决方案以在这片新兴领域抢占先机。根据 Semi analysis，共封装光学供应链中的关键设备供应商包括是德科技、Ficontec、泰瑞达、爱德万测试、FormFactor、致茂电子、安立与 Multilane。是德科技是该领域的重要参与者，这家企业以提供高品质高速测试设备而闻名，例如近期针对 1.6T 光模块测试推出的两款新型示波器，凭借其市场地位，该公司有望充分受益于共封装光学技术的发展趋势。Ficontec 在光子测试领域拥有坚实基础，现正积极拓展电气测试能力，其核心优势之一是晶圆级光子测试技术，这项技术显著提升了传统低效光子测试流程的效率，该公司近期推出了新型光子集成电路高通量晶圆级测试设备，兼容现有半导体自动化测试架构，号称业界首款双面晶圆测试机。泰瑞达同样是该领域的重要参与者，在电气测试领域积淀深厚，该公司近期收购了一家专注封装光测试的初创企业，彰显了其构建 CPO 测试能力的战略决心。致茂电子一直为 3D 传感与光通信领域的激光二极管提供光子测试设备，具备利用此项专长进军 CPO 领域的潜力，但该公司当前的创新步伐略落后于部分竞争对手。

综上所述，CPO 交换机供应链呈现出高度专业化与垂直整合并存的特征，从光子芯片制造、先进封装到系统测试，各环节的技术壁垒与价值分布正在快速演化。随着英伟达、博通等龙头厂商加速 CPO 产品量产节奏，供应链上游的光学元件、封装测试与设备厂商将迎来重要的成长机遇，而具备技术先发优势与客户粘性的企业有望在竞争中脱颖而出。

4. 投资建议

英伟达 NVLink SerDes 速率已从 56Gbps 迭代至 224Gbps，支撑单芯片互联带宽跨越式升级，但 224G 以上信号高频衰减剧增且 SerDes 功耗占比随速率攀升，倒逼算力互联介质从“电”向“光”加速演进。在机柜内部，即将推出的 Rubin Ultra 以 144 颗 GPU 构建 1.5PB/s PB 级 Scale-up 网络，其 4-Canister 双层架构直接回应上述挑战：第一层正交背板因 224G SerDes 严苛的信号完整性要求必须采用 M9 级超低损耗 CCL 材料；第二

层跨 Canister 互联则通过 72 颗 NVSwitch 与 648 颗 3.2T NPO 光引擎实现光电近封装，GPU 与光引擎配比高达 1:4.5，标志着“光入柜内”趋势确立。在机柜间的 Scale-out 网络侧，英伟达 CPO 交换机产品矩阵已抢先卡位——Quantum X3450 作为全球首款量产 CPO 交换机以 115.2T 带宽与可拆卸光引擎设计树立标杆，Spectrum X 平台（6810/6800）更以 102.4T-409.6T 梯度化带宽覆盖以太网生态，量产后有望完善从 InfiniBand 到以太网的全场景布局，形成机柜内 NPO 与网络侧 CPO 协同的技术闭环。

基于上述技术演进路径，我们认为 PCB 覆铜板向 M9 等级升级、光电共封装（CPO）及“光入柜内”已成为 2026-2027 年算力基础设施的确定性技术趋势。随着 Rubin Ultra 机柜落地及 CPO 交换机规模化放量，M9 级覆铜板、NPO/CPO 光引擎、外部激光源、光纤连接单元等核心零部件需求将迎来爆发式增长。建议 2026 年重点关注两大投资主线：一是 PCB M9 材料产业链，二是 CPO 及光入柜内所对应的光芯片、光器件、光引擎等光互联产业链。具备供应链卡位优势的国内龙头厂商有望在技术迭代与需求扩容的双重驱动下率先受益，迎来价值重估机遇，建议关注：

M9 PCB 产业链：菲利华、东材科技、生益科技、胜宏科技、沪电股份、深南电路、东山精密等

CPO 产业链：致尚科技、长光华芯、源杰科技、仕佳光子、太辰光、炬光科技、罗博特科等

5. 风险提示

算力互联需求不及预期：下游 AI 算力建设投入力度、算力网络带宽扩容规模若未达市场预期，将直接导致客户对网络互联相关产品的采购需求放缓，进而对相关公司的营收增长及盈利表现产生不利影响。

客户拓展及份额提升不及预期：相关企业若未能按预期开拓潜在客户资源，或在现有核心客户供应链中的份额提升进度低于预期，将制约其市场渗透率提升，进而影响业绩增长的持续性。

产品研发及量产落地不及预期：针对高潜力的算力应用场景产品，若相关公司在技术研发突破、产品迭代升级或规模化量产落地环节未达预期，将错失市场发展机遇，对公司长期业绩增长构成制约。

行业竞争加剧：随着新兴 AI 基础设施产品商业化进程加快，行业参与者增多可能导致竞争持续加剧，若相关公司未能维持技术壁垒、成本控制或客户服务优势，其市场份额存在下滑风险，进而影响盈利能力。

免责声明

东吴证券股份有限公司经中国证券监督管理委员会批准，已具备证券投资咨询业务资格。

本研究报告仅供东吴证券股份有限公司（以下简称“本公司”）的客户使用。本公司不会因接收人收到本报告而视其为客户。在任何情况下，本报告中的信息或所表述的意见并不构成对任何人的投资建议，本公司及作者不对任何人因使用本报告中的内容所导致的任何后果负任何责任。任何形式的分享证券投资收益或者分担证券投资损失的书面或口头承诺均为无效。

在法律许可的情况下，东吴证券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并进行交易，还可能为这些公司提供投资银行服务或其他服务。

市场有风险，投资需谨慎。本报告是基于本公司分析师认为可靠且已公开的信息，本公司力求但不保证这些信息的准确性和完整性，也不保证文中观点或陈述不会发生任何变更，在不同时期，本公司可发出与本报告所载资料、意见及推测不一致的报告。

本报告的版权归本公司所有，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。经授权刊载、转发本报告或者摘要的，应当注明出处为东吴证券研究所，并注明本报告发布人和发布日期，提示使用本报告的风险，且不得对本报告进行有悖原意的引用、删节和修改。未经授权或未按要求刊载、转发本报告的，应当承担相应的法律责任。本公司将保留向其追究法律责任的权利。

东吴证券投资评级标准

投资评级基于分析师对报告发布日后 6 至 12 个月内行业或公司回报潜力相对基准表现的预期（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，新三板基准指数为三板成指（针对协议转让标的）或三板做市指数（针对做市转让标的），北交所基准指数为北证 50 指数），具体如下：

公司投资评级：

买入：预期未来 6 个月个股涨跌幅相对基准在 15%以上；

增持：预期未来 6 个月个股涨跌幅相对基准介于 5%与 15%之间；

中性：预期未来 6 个月个股涨跌幅相对基准介于-5%与 5%之间；

减持：预期未来 6 个月个股涨跌幅相对基准介于-15%与-5%之间；

卖出：预期未来 6 个月个股涨跌幅相对基准在-15%以下。

行业投资评级：

增持：预期未来 6 个月内，行业指数相对强于基准 5%以上；

中性：预期未来 6 个月内，行业指数相对基准-5%与 5%；

减持：预期未来 6 个月内，行业指数相对弱于基准 5%以上。

我们在此提醒您，不同证券研究机构采用不同的评级术语及评级标准。我们采用的是相对评级体系，表示投资的相对比重建议。投资者买入或者卖出证券的决定应当充分考虑自身特定状况，如具体投资目的、财务状况以及特定需求等，并完整理解和使用本报告内容，不应视本报告为做出投资决策的唯一因素。

东吴证券研究所
苏州工业园区星阳街 5 号

邮政编码：215021

传真：（0512）62938527

公司网址：<http://www.dwzq.com.cn>