

电子行业深度报告

超节点方案重塑 AI 算力产业趋势, 开启系统级算力新周期

增持（维持）

2026 年 08 月 27 日

证券分析师 陈海进

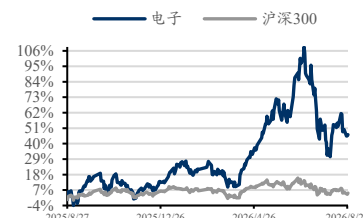
执业证书: S0600525020001

chenhj@dwzq.com.cn

投资要点

- **算力供需共振下, 超节点成为 AI 集群演进的必然路径。**全球 AI 前沿模型能力正从单轮问答走向复杂任务处理, 长上下文、复杂推理与 Agent 化应用持续拉长推理链条, 推动训练与推理算力需求上行。与此同时, 海外云厂商加速扩建 AI 基础设施, 国内大模型应用放量和国产算力供给释放形成共振, 算力基础设施的主要矛盾正在从“扩充算力卡”走向“提升有效算力”。超节点通过扩大高带宽协同域、缩短高频通信路径, 将分散算力组织为可调度、可扩展的系统算力, 成为突破传统集群效率瓶颈、实现大规模协同的核心破局点。
- **超节点架构升级的核心变量: 高带宽互联与整柜工程能力。**(一) 通信层面, Scale Up 负责在超节点内部承接高频数据交换, Scale Out 负责多柜、多集群的横向扩展, 二者边界正从物理服务器边界转向性能边界; 英伟达 NVLink 闭环体系、华为灵衢 UB 国产体系、AMD/UALink 海外开放路线以及阿里、字节、腾讯等云厂商开放互联方案并行演进, 共同推动超节点从架构探索走向产品化。(二) 工程层面, 百千瓦级整柜功耗、高密度连接和长周期稳定运行, 使散热、供电、结构和运维从单点设计升级为整柜系统工程。超节点竞争由此从单卡性能延伸到整柜协同能力, 价值量也从计算卡外溢至互联交换、液冷供电、系统验证和整柜交付环节。
- **超节点重塑服务器厂商价值链, 整柜交付与客户生态成分化关键。**超节点推动 AI 服务器需求从单机采购转向整柜系统和智算中心级方案交付, 服务器厂商的价值边界也从整机制造扩展到整柜集成、系统测试和现场部署。工业富联受益于全球头部客户整柜制造和高速网络设备放量; 华勤技术依托平台型 ODM 能力向 AI 服务器、超节点和数据中心产品延伸; 浪潮信息、紫光股份、超聚变和宁畅则分别围绕国产 AI 服务器、网络底座、运营商液冷整柜和高密定制化场景形成差异化卡位。随着超节点从样机验证走向规模交付, 客户绑定深度、整柜工程能力和项目复制效率将成为服务器厂商业绩兑现的核心支撑。
- **风险提示:** 算力需求不及预期; 技术路线与商业化落地; 整柜工程量交付不及预期; 行业竞争加剧与盈利水平下滑。

行业走势



相关研究

《海外算力周跟踪: 关注 CPO 商业化落地进程》

2026-08-23

《端侧 AI 周跟踪: Agent 模型加速落地, 端云协同提升端侧 AI 价值》

2026-08-16

内容目录

1. 全球算力需求激增：AI 应用升级与基础设施扩张形成共振.....	7
1.1. 前沿模型与 Agent 化应用升级推动算力需求持续攀升.....	7
1.2. 海外厂商持续扩建 AI 基础设施，算力供给进入加速释放阶段.....	7
2. 国产算力需求：国内 AI 应用放量推动国产算力需求释放.....	8
3. 算力集群演进：算力集群形态与算力需求演进高度同步演化.....	9
4. 超节点：国产算力需求的破局之策.....	9
4.1. 超节点架构加速规模化落地，推动国产 AI 集群向高密度协同演进.....	9
4.2. 高密度集成与高速互联构筑超节点核心能力底座.....	10
5. 超节点系统架构升级：高带宽互联与整柜工程成为核心变量.....	11
5.1. 通信是超节点的第一主线：Scale Up 与 Scale Out.....	11
5.1.1. Scale Up：高带宽域内互联决定系统有效算力.....	11
5.1.2. Scale Up 路线分化：专有互联、总线化互联、云侧自定义与开放以太网并行演进.....	13
5.1.3. Scale Out：超节点之间的横向扩展决定集群效率.....	14
5.1.4. Scale Up 与 Scale Out 边界变化：从服务器边界转向性能边界.....	15
5.1.5. 光互联与 CPO：高带宽域扩大后的下一阶段通信增量.....	16
5.2. 工程系统是超节点落地的第二主线：整柜级交付成为产品化门槛.....	17
5.2.1. 散热：从服务器级散热走向整柜级液冷.....	18
5.2.2. 供电：从电源模块走向整柜级高功率供电.....	18
5.2.3. 结构与线缆：高密度系统中的机械工程问题.....	19
5.2.4. 可靠性：从单机可用性走向系统级容错.....	20
5.2.5. 软件与运维：硬件互联决定上限，系统软件决定释放效率.....	21
6. 主流超节点产品演进：闭环生态、国产体系与开放路线并行.....	22
6.1. 英伟达：从 NVLink 整柜系统走向 AI Factory.....	22
6.1.1. GB200 NVL72 与 GB300 NVL72：整柜级超节点形态确立并向推理增强.....	23
6.1.2. Vera Rubin、Rubin Ultra 与 Kyber Rack：从单柜 NVLink 域走向 POD 级 AI Factory.....	27
6.2. 海外开放路线：AMD Helios 与 UALink 生态形成补充.....	33
6.2.1. AMD Helios：海外非英伟达生态进入 rack-scale 系统竞争.....	34
6.2.2. UALink：开放 Scale Up 标准为非闭环生态提供接口.....	35
6.3. 华为：以系统级扩展构建国产超节点体系.....	36
6.3.1. Atlas 900 A3 / 昇腾 384：国产超节点起点.....	38
6.3.2. Atlas 950/960 SuperPoD：国产 Scale Up 域持续扩大.....	40
6.4. 中科曙光：开放架构下的高密度超节点系统.....	41
6.4.1. scaleX640：单机柜级高密度超节点.....	42
6.4.2. scaleFabric 与 scaleX 万卡超集群：从单柜超节点走向集群级基础设施.....	43
6.5. 云厂商与开放互联：真实负载牵引架构创新.....	44
6.5.1. 阿里磐久 AL128：云侧资源池化牵引超节点重构.....	44
6.5.2. 字节 EthLink 与腾讯 ETH-X：以太网 Scale Up 的协议探索.....	45
6.5.3. ETH+、ERack+ 与 UPN 512：国内开放超节点从协议走向样机.....	45
7. 超节点重塑服务器厂商价值链：整柜交付与客户生态成为分化核心.....	47

- 7.1. 超节点推动服务器厂商角色升级：从整机制造走向整柜系统交付.....47
- 7.2. 工业富联：全球 AI 服务器与整柜制造链核心平台48
- 7.3. 华勤技术：平台型 ODM 向 AI 服务器和数据中心产品延伸50
- 7.4. 浪潮信息：国内 AI 服务器龙头，液冷整柜与智算中心交付能力突出52
- 7.5. 紫光股份（新华三）：网络底座与政企智算优势强化系统交付能力.....53
- 7.6. 超聚变：运营商与政企场景中的整机柜液冷系统交付平台.....56
- 7.7. 宁畅：以液冷整柜与高密 AI 服务器切入智算场景，超节点卡位仍待规模验证58
- 8. 系统仿真与设计验证：从工程复杂度走向可复制交付.....60
 - 8.1. 单柜功耗跃迁推动设计验证前置.....60
 - 8.2. 热管理仿真与虚拟调试决定液冷方案收敛效率.....61
 - 8.3. 数字孪生与多物理场仿真推动整柜设计从经验走向模型驱动.....62
 - 8.4. 整柜测试与交付流程决定规模复制能力.....63
- 9. 风险提示64

图表目录

图 1:	2021-2027E 海外 CSP 资本开支 (亿美元)	8
图 2:	2025 年国产算力卡出货量	9
图 3:	中科曙光 scaleX 万卡级超集群首次亮相	10
图 4:	基础计算架构通信范式的演变	11
图 5:	大模型采取并行策略训练	12
图 6:	超节点优化推理架构	12
图 7:	GB200 NVL72 以整柜级 NVLink 互联证明高带宽系统交付成为新壁垒	13
图 8:	Scale up & Scale out 组网示意图	14
图 9:	AI 流量在集群中表现的负载不均和拥塞问题	15
图 10:	Scale-Up 和 Scale-Out 融合和独立组网对比	16
图 11:	英伟达 GB200 内部使用铜连接	16
图 12:	Scale-up 互联由铜互联向光互联迁移	17
图 13:	CPO 有望率先在跨机柜 Scale-up 场景落地	17
图 14:	整柜系统集成覆盖液冷、电力、线缆与结构工程	17
图 15:	GB200 液冷计算托盘内部结构	18
图 16:	机柜背板通过 manifold 供水	18
图 17:	机柜功耗持续攀升带动 Power shelf 和 PSU 需求	19
图 18:	采用 Busbar 为各节点提供电源	19
图 19:	OEX 与 Cable Tray 方案对比	19
图 20:	超节点系统级 RAS 管理架构	20
图 21:	硬件可靠性策略	20
图 22:	系统软件支撑集群调度、监控与恢复	21
图 23:	NVLink 的 Scale up 域迅速扩张	22
图 24:	GB200 NVL72 由计算托盘走向整柜级 Scale Up 系统	23
图 25:	英伟达首款 CPU-GPU 超级芯片	24
图 26:	DGX GH200 超级计算机通过 NVLink 连接最多 256 个 GH200 超级芯片	24
图 27:	GH200 NVL32 单柜级 NVLink 系统结构	25
图 28:	Blackwell NVL72 推理吞吐与能效较 Hopper 显著提升	25
图 29:	GB200 NVL72 以整柜级模块化布局确立 72 GPU 标准交付形态	26
图 30:	GB200 NVL72 以 2 CPU + 4 GPU 计算托盘为基本单元	26
图 31:	GB300 NVL72 网络带宽较 GB200 NVL72 提升 2 倍, FP4 推理性能提升 1.5 倍	27
图 32:	Rubin NVL72 推理性能与能效进一步提升	28
图 33:	Rubin NVL72 以更少 GPU 完成同等训练任务	28
图 34:	Rubin 推理成本降至 Blackwell 的十分之一	28
图 35:	Vera Rubin NVFP4 推理性能达 3.6EFLOPS, 较 Blackwell 提升 5 倍	29
图 36:	Vera Rubin 进一步强化 CPU-GPU 一致性内存与 NVLink-C2C 互联	29
图 37:	POD 以五类机架系统将 NVL72 扩展为多机架 AI Factory 单元	30
图 38:	Vera Rubin POD 由 40 个机架、1,152 个 Rubin GPU 组成	30
图 39:	Vera Rubin 通过“七芯片、五机架”协同, 将 AI Factory 竞争推进到系统级 token 产出效率	31
图 40:	Rubin Ultra NVL576 系统性能约为 GB300 NVL72 的 14 倍	32
图 41:	英伟达 Polyphe 原型机, 基于 GB200 的多机架 NVL576 架构	32

图 42:	Kyber 是下一代 MGX 机架, 将首先应用于 Rubin Ultra	33
图 43:	NVLink Fusion 显示英伟达正将 NVLink 闭环能力向开放芯片生态延伸	33
图 44:	UALink 面向 Scale Up 机架基础设施构建多节点、多平面交换网络	34
图 45:	AMD Helios 以开放机架切入超节点系统竞争	35
图 46:	Helios 计算托盘将 CPU、GPU、DPU 与 AI NIC 集成在同一模块	35
图 47:	UALink 分层架构	35
图 48:	UALink 联盟董事会成员	35
图 49:	UALink 演进时间线	36
图 50:	灵衢 UB 通过两级 scale up 交换组织 384 NPU + 192 CPU 协同计算域	37
图 51:	灵衢协议栈	37
图 52:	超节点集群组网架构 (以昇腾 384 超节点为例)	38
图 53:	昇腾 384 将国产算力从服务器堆叠推进到多柜超节点形态	39
图 54:	Atlas 900 A3 以 384 NPU 高速互联、统一编址和百纳秒级时延构建国产 Scale Up 能力	39
图 55:	Atlas 900 A3 机柜整机示意图	39
图 56:	Atlas 900 A3 部署落地情况 (截至 2025 年 9 月 18 日)	39
图 57:	昇腾芯片路线持续提升算力、互联带宽与内存	40
图 58:	Atlas 950 SuperPoD 将采用昇腾 950DT, Scale Up 域扩展至 8,192 NPU	40
图 59:	Atlas 960 SuperPoD 进一步扩大至 15,488 NPU	40
图 60:	Atlas 950 SuperPoD 具备 16.3PB/s 互联带宽, 推理吞吐达 19.6mn TPS	41
图 61:	scaleFabric 以软硬件开放适配承接多元国产算力生态	42
图 62:	scaleX640 以 640 卡高密整柜实现开放架构下的单柜超节点产品化	43
图 63:	scaleFabric 构建从网卡、交换机到管理软件的集群级高速互联底座	43
图 64:	scaleX 万卡超集群标志中科曙光从单柜超节点走向集群级基础设施	44
图 65:	scaleX 万卡超集群由 16 个 scaleX640 互联组成, 形成 10240 卡、超 5EFLOPS 系统能力	44
图 66:	字节 EthLink 协议栈	45
图 67:	ETH-X Scale U 互联协议示意图	45
图 68:	ETH+ Scale up 协议架构	46
图 69:	基于 ETH+ 光互联的超节点解决方案	46
图 70:	从 L10 到 L11 整机柜、L12 多机柜交付	47
图 71:	2025Q4 分架构营收 (亿美元), 非 x86 狂飙	48
图 72:	2025Q4 全球厂商份额, ODM 首次超过 OEM 总和	48
图 73:	2026 一季度 AI GPU 机柜、AI ASIC 服务器与 800G 交换机高增	50
图 74:	“3+N+3” 产品平台覆盖多元终端与服务器, 平台型 ODM 能力向数据中心延伸	50
图 75:	数据中心业务持续高增, AI 服务器与超节点加速产品升级	51
图 76:	ETH-X 原型机在华勤技术东莞智能制造基地点亮	52
图 77:	元脑 SD200 与 HC1000 分别强化国产芯片互联和大规模推理能力	53
图 78:	新华三 102.4T 智算交换机通过光电并行提升密度、吞吐与低时延能力	54
图 79:	新华三智算广域网融合核心路由器与 800G 交换机, 支撑千公里级无损传输	54

图 80: UniPoD S80000 支持 1024 卡 Scale Up 互联与 128 卡单柜部署	55
图 81: 全栈协同定义新一代 AI 基础设施	56
图 82: 超聚变产品体系覆盖服务器、操作系统与 AI 解决方案, 强化国产智算基础设施交付能力	56
图 83: FusionPoD for AI 集成液冷、供电与运维管理, 支撑 64—128 卡整机柜交付	57
图 84: 静音全液冷一体化系统联合创新样板点	58
图 85: 宁畅 B8000 以 42U 整机柜形态整合计算、交换、供电与液冷单元	59
图 86: DeepSeek 大模型一体机解决方案产品矩阵	60
图 87: 宁畅“全栈全液”AI 基础设施方案架构	60
图 88: AI 整柜功耗从十千瓦级跃升至百千瓦级, 并向兆瓦级系统演进	61
图 89: 数据中心温湿度与冷冻水工况仿真	62
图 90: 冷却系统虚拟调试与控制策略验证	62
图 91: 数字孪生平台贯通设计、运行、资产与仿真	62
图 92: AI Factory 多物理场仿真体系	62
图 93: 机架集成流程覆盖设计、装配、烧机、交付与现场部署	63
表 1: 并行策略下的 TP、EP 对带宽有极高要求	12
表 2: Scale up 协议全景	13

1. 全球算力需求激增：AI 应用升级与基础设施扩张形成共振

1.1. 前沿模型与 Agent 化应用升级推动算力需求持续攀升

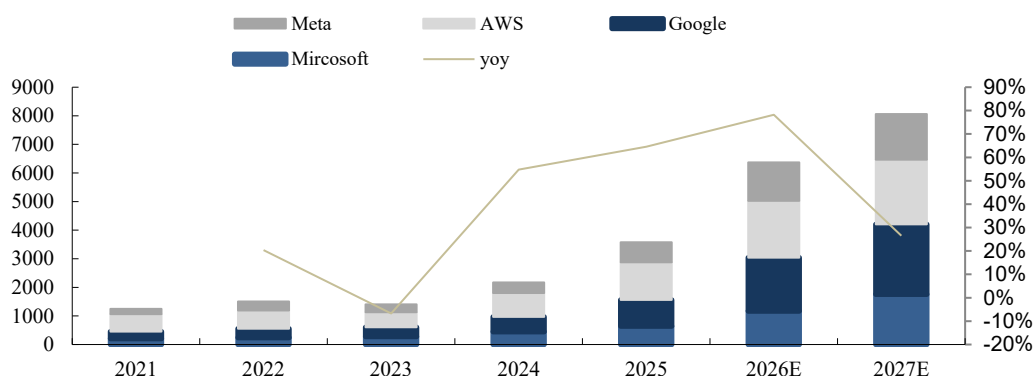
前沿模型迭代与应用升级正推动 AI 工作负载加重,全球算力需求持续攀升。近期, GPT-5.5、DeepSeek V4 几乎同期上线,突破百万 Token 长上下文能力,可处理多步骤推理、长文档理解及多轮交互等复杂任务,使单次推理可覆盖数百页文档信息。这些功能升级背后,需要大量训练算力的投入,成为推动算力需求持续增长的重要因素。当前模型升级的趋势已从“更会回答”转向“更会处理复杂任务”,算力消耗激增不仅源于模型参数扩张,更受推理阶段任务复杂度影响。单任务计算量、显存占用及数据传输均显著增加,从而提升了整体算力负载。Token 是大模型处理语言的基本单位,模型所有输入与输出均以 Token 计量,其数量直接决定计算量——上下文越长、任务越复杂,Token 消耗越高。援引高盛预测数据,随着消费端和企业端 AI Agent 加速普及,全球 Token 消耗量将从 2026 年的约 5 千万亿 token/月增长至 2030 年的 120 千万亿 token/月,四年增长约 24 倍, CAGR 约 121%。

AI Agent 应用快速扩张和多模态任务激增,推动推理算力需求激增成为算力需求增长的核心引擎。 Agent 具有任务拆解、工具调用、多轮反思与自主执行等特性,让单次指令的推理链条变长、调用层数增加、上下文持续累积。结合多模态交互能力,单个任务往往需要跨模态信息融合与复杂推理,导致 Token 调用量呈指数级增长,形成“用量乘数效应”,进一步放大推理算力需求。据 IDC 预测,2027 年推理算力在智能算力总需求中的占比将超过 70%,推理算力需求将接棒成为主引擎。

1.2. 海外厂商持续扩建 AI 基础设施,算力供给进入加速释放阶段

海外云厂商加码资本开支,验证算力需求持续性。算力需求的爆发式增长远超市场预期,但当前算力供给存在短期瓶颈,受硬件、架构及资源调度等多重因素制约,短期内难以快速匹配需求增长。海外头部云厂商正在加速算力供给以应对需求压力,资本投入持续上升成为直接验证。2025 年,Alphabet 关于服务器、网络设备和数据中心等技术基础设施支出 914 亿美元,公司预计 2026 年将增至 1800 亿美元左右; Meta 将 2026 年资本开支指引由 1,150 亿-1,350 亿美元上调至 1,300 亿-1,450 亿美元,多家 hyperscaler 几乎同步加码 AI 服务器与数据中心基础设施,表明全球 AI 正由技术验证加速迈向生产级部署,算力需求仍处于高景气上行区间,进一步支撑算力产业的持续增长预期。

图1：2021-2027E 海外 CSP 资本开支（亿美元）



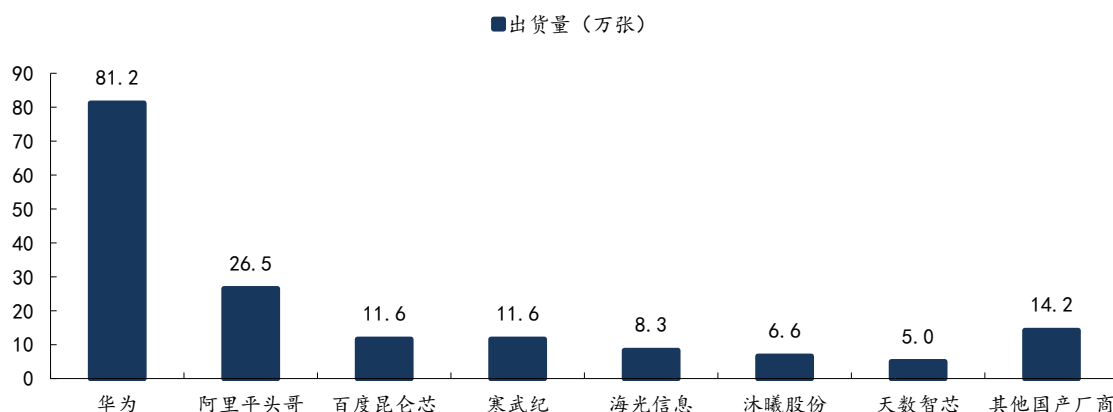
数据来源：Bloomberg 及其预测，东吴证券研究所

2. 国产算力需求：国内 AI 应用放量推动国产算力需求释放

国产算力需求与基础设施建设同步提速，产业进入加速发展阶段。一方面，以 DeepSeek V4 为例的一众国产大模型不断升级，加速进入商业化规模化应用，带动训练与推理负载同步提升。据国家统计局数据，到今年3月份，日均 Token 调用量突破 140 万亿，比上年末增长超 40%。另一方面，虽然当前全球高端 AI 算力仍高度依赖 NVIDIA 等海外厂商，但 AI 芯片同时受到出口管制、供应链限制及地缘政策等因素影响，海外算力供给存在不确定性，促使国内算力产业在高端芯片、服务器及自主生态建设方面具备更强的国产替代需求，也进一步推动本土 AI 算力体系加速发展。在需求扩张与自主可控趋势推动下，国内算力基础设施已进入加速建设阶段，智能算力规模持续扩张。国家数据局披露，截至 2026 年3月底，全国智能算力总规模已达 188 万 PFLOPS(FP16)，其中八大国家算力枢纽节点占比超过 80%。“东数西算”等全国一体化算力网络建设也在持续推进，通过跨区域资源调度优化算力布局，缓解区域性供需错配问题。且根据 IDC 预测，2026-2028 年中国的智能算力规模仍将维持约 40%的年均增速，智能算力在整体算力市场中的占比有望进一步提升，中国 AI 算力产业增长空间仍具备较强释放潜力。

国产算力需求的主要矛盾正在由“缺卡”转向“高效用卡”，系统组织能力成为关键。算力卡，即 AI 加速卡，以 GPU、ASIC 及相关异构芯片为主要形态，是承载大模型训练与推理任务的核心硬件载体，通过提供高并行计算能力，支撑大规模参数模型的迭代与部署。根据 IDC 的数据显示，2025 年中国市场的 AI 加速卡出货量约为 400 万张，其中国产算力卡出货量约为 165 万张，市场份额约为 41%。国产加速卡供给缺口逐步缓解，但应用端仍存突出问题：单卡层面，据中国电子技术标准化研究院测试显示，国产 GPU 实际可用算力约为标称值的 50-70%，还面临着功耗大与生态成熟度不足等问题。单纯粗放堆卡已无法匹配复杂算力需求，需转向系统级优化。

图2：2025 年国产算力卡出货量



数据来源：IDC，通信世界，东吴证券研究所

3. 算力集群演进：算力集群形态与算力需求演进高度同步演化

随着大模型计算复杂度持续提升，传统算力集群的效率瓶颈逐步显现。在通用计算阶段，算力供给以 CPU 为核心，节点间协同程度有限，主要依赖单机或小规模集群满足基础计算需求。随着深度学习模型规模扩张及 Transformer 架构普及，GPU 逐步成为算力核心，系统演进为多 GPU 协同的 Scale Out 集群架构。在该架构下，GPU 通过 NVLink 实现机内高速互联，跨节点则依赖 InfiniBand 等网络协议进行通信。然而，随着千亿级甚至万亿级模型训练普及，以及张量并行、专家并行与序列并行等多维并行方式的广泛应用，跨节点通信频率与数据规模显著上升，大规模梯度同步和 All-to-All 通信开销持续上升，通信延迟与带宽拥塞逐步成为制约有效算力释放的核心瓶颈，并进一步向功耗、散热及运维复杂度等系统层面传导。传统横向扩展模式已难以匹配高频通信与大规模并行计算需求，算力集群开始从以服务器为核心的松耦合架构，向更高密度、更强协同的一体化计算单元演进。

4. 超节点：国产算力需求的破局之策

4.1. 超节点架构加速规模化落地，推动国产 AI 集群向高密度协同演进

超节点架构正成为国产 AI 算力突破通信瓶颈、促进有效算力释放的核心演进方向。超节点的核心定义是以高带宽域为边界的大规模 Scale Up 计算单元，通过高速互联协议、交换节点、统一寻址访存机制及系统调度，将数十至数百个计算模组组织成低时延、高带宽、可编排的逻辑计算体。其核心价值并非单纯增加卡数，而是缩短高频通信路径、降低同步开销，实现峰值算力向有效吞吐的转化。超节点架构在产业端形成规模化落地，成为新一代国产智算集群的发展重心，国内厂商积极布局：华为昇腾 384、阿里磐久 128、百度天池 512 等超节点架构均实现“百卡级”超节点架构突破。2025 年 12 月，中科曙光 scaleX 万卡超集群首次亮相。2026 年 2 月 5 日，3 套 scaleX 万卡超集群系统在国家

超算互联网郑州核心节点同步上线,建成全国首个部署3万张国产AI加速卡的算力池。在2026数字中国建设峰会上,华为称昇腾384已部署超500套,实现规模化商用,标志着国产AI基础设施正从传统分布式集群向高密度系统级算力架构加速演进。

图3: 中科曙光 scaleX 万卡级超集群首次亮相



数据来源: 中科曙光公司官网, 东吴证券研究所

超节点打破传统服务器物理边界, 重构产业协同创新的算力新范式。超节点作为系统级创新架构, 通过极致的内部高速互联与全栈协同设计, 从底层架构破局, 实现大规模资源紧密耦合与高效协同。有效解决传统服务器集群在大模型训练中面临的通信墙、功耗墙和复杂度墙三大技术瓶颈: 传统集群通常采用以太网或 PCIe 组网, 供电与散热相对分散, 计算节点呈松散堆叠形态, 导致延迟高、带宽不足、散热受限、部署运维复杂, 超节点则通过柜内高速互联、原生液冷、集中供电与统一管理, 大幅提升通信效率、降低能耗、简化调度。超节点架构能通过物理与逻辑层面的深度紧耦合, 实现算力集群的效率跃升, 从而支撑并行计算任务, 加速 GPU 间的参数交换和数据同步, 重构 AI 算力基础设施的基本范式。

4.2. 高密度集成与高速互联构筑超节点核心能力底座

高密度集成构筑硬件根基, 系统级融合升维助力超节点架构的切实落地。超节点的核心是在极高空间密度下集成海量计算的能力, 高密度集成设计能使其在有限物理空间内实现算力的规模化, 能够显著提升算力密度。硬件层面, 采用整机柜设计, 在机柜内部署高速的电或光背板, 取代传统的外部线缆。通过供电、散热、计算、交换模块一体化设计, 提升部署密度和效率。超节点高算力背后的巨大功耗促使工程升级, 以实现能效利用率优化。供电方面, 超节点数据中心从交流供电向高压直流供电转型, 提升了从电网到芯片的端到端供电效率。散热方面, 冷板式液冷成为标配技术, 有效实现系统的

热量传导。

超节点依托高密度集成、高速互联、全局协同核心优势，有望成为未来智算中心核心部署形态。与传统的计算节点相比，超节点可通过特有互联协议实现 AI 加速卡的高速互联，构建高带宽域优势。英伟达的 NVLink 6.0 每 GPU 带宽已提升至 3.6TB/s。在时延方面，Scale Out 采用的 IB 和 RoCEv2 的通信网络技术，通信时延高达 10 微秒。而超节点的 Scale Up 架构构建机柜内总线通信，不需要经过跨服务器通信的复杂路径，通信时延可降至百纳秒级别。超节点的高速互联和协同所带来的高带宽与低时延优势，将打破堆服务器扩充算力的旧模式。且随着千卡级超节点全面进入常态化部署阶段，超节点规模有望进一步扩大，百万卡的超大规模算力集群将重塑 AI 算力和硬件市场新格局。

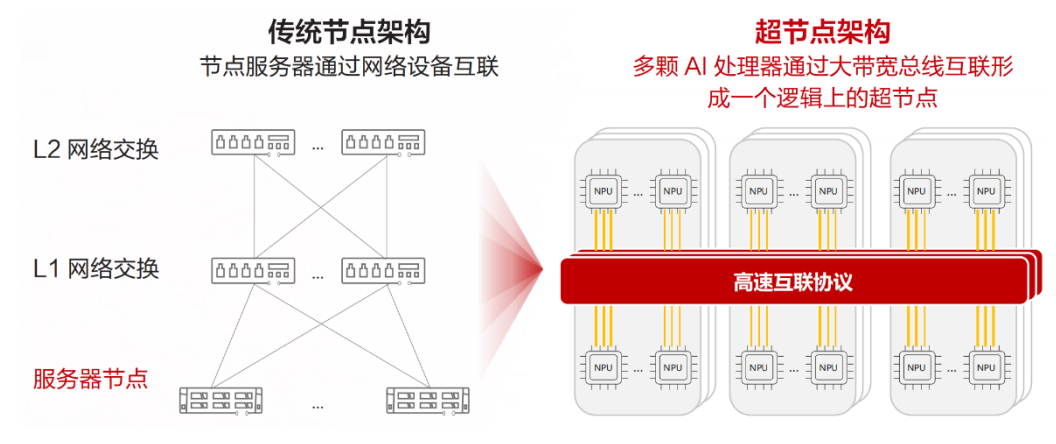
5. 超节点系统架构升级：高带宽互联与整柜工程成为核心变量

5.1. 通信是超节点的第一主线：Scale Up 与 Scale Out

通信架构决定超节点能否把分散计算资源转化为系统算力。大模型训练中的张量并行、专家并行和集合通信，以及推理侧长上下文、KV Cache 和高并发服务，都会放大计算模组之间的数据交换需求。超节点的关键变化，是通过更大的高带宽域（HBD）将高频通信收敛到受控范围内，降低跨服务器通信带来的等待和同步损耗。

通信系统升级正在推动服务器价值量从单机整机向互联环节迁移。过去 AI 服务器竞争更多围绕加速卡、主板和单机散热展开；进入超节点阶段后，计算托盘、Switch Tray、高速铜缆、背板/中板、连接器、光互联和通信软件共同影响整柜算力释放。通信能力由此从配套环节上升为产品形态和工程壁垒的核心变量。

图4：基础计算架构通信范式的演变

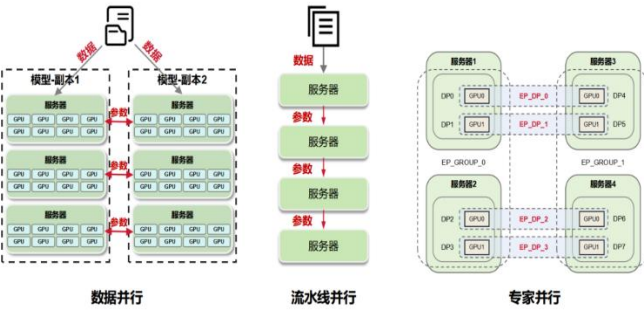


数据来源：《超节点发展报告》，东吴证券研究所

5.1.1. Scale Up: 高带宽域内互联决定系统有效算力

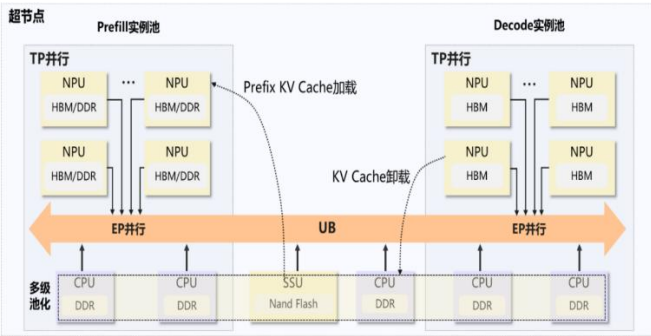
Scale Up 的核心价值，是在超节点内部构建高带宽、低时延的协同计算域。训练侧 TP、EP 对应的 All-Reduce、All-to-All 通信，对带宽、时延和同步效率要求较高；推理侧长上下文和 KV Cache 也会增加跨卡访问需求。通信路径越长，计算单元等待时间越明显，峰值算力向有效吞吐的转化效率越低。

图5：大模型采取并行策略训练



数据来源：鲜枣课堂，东吴证券研究所

图6：超节点优化推理架构



数据来源：《基于灵衢的超节点参考架构白皮书》，东吴证券研究所

表1：并行策略下的 TP、EP 对带宽有极高要求

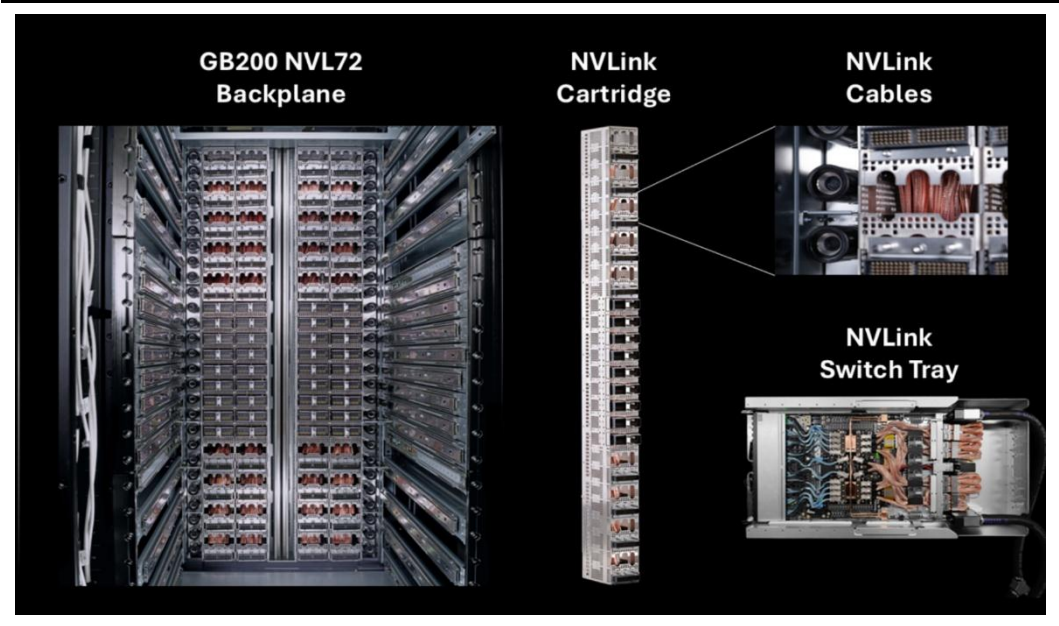
并行模式	主要通信方式	主要发生位置	对网络的主要要求
张量并行 (TP)	All-Reduce	单机或单超节点内	极高带宽、低时延、同步性
专家并行 (EP)	All-To-All	超节点内优先	极高带宽、低尾延迟、动态调度
流水并行 (PP)	Send/Recv	多节点之间	稳定点到点传输
数据并行 (DP)	All-Reduce	全局集群范围	大规模扩展能力与全局利用率

数据来源：《超节点技术体系白皮书》，鲜枣课堂，东吴证券研究所

从单机八卡走向整柜级超节点后，Scale Up 的实现难度明显提升。NVLink/NVSwitch 通过交换式互联把多颗 GPU 组织到同一高带宽域内，华为 UB、海光 HSL 等路线则通过统一编址、内存语义和多层互联扩大国产算力的协同边界。技术路线不同，但目标一致：让更多计算单元在逻辑上更接近一个统一计算体。

高带宽域扩大后，系统竞争从单卡性能延伸到整柜协同能力。加速卡之间能否稳定、低损耗地完成数据交换，开始与单卡算力共同决定产品竞争力。交换节点、高速连接、板级设计、通信库和整柜验证因此获得更高价值位置，服务器厂商的能力边界也从整机装配延伸到高带宽系统交付。

图7：GB200 NVL72 以整柜级 NVLink 互联证明高带宽系统交付成为新壁垒



数据来源：英伟达官网，东吴证券研究所

5.1.2. Scale Up 路线分化：专有互联、总线化互联、云侧自定义与开放以太网并行演进

Scale Up 路线分化，本质上是性能确定性、生态控制权和开放兼容性之间的不同取舍。英伟达 NVLink/NVSwitch 代表闭环高性能路线，通过 GPU、互联协议、交换芯片和软件生态同步迭代，率先形成成熟的整柜产品形态。AMD Infinity Fabric 同样属于厂商自有互联体系，更强调芯片与平台内部协同。此类路线的优势在于控制力强、产品节奏清晰，短期仍是高端超节点的主导形态。

表2：Scale up 协议全景

Scale up 协议	主导方	物理层	组网规模（卡）	超节点产品
NVLink	NVIDIA	私有 SerDes	单层，≤576	NVL72/576
UALink	AMD/Intel/Google 等	以太网 PHY 或 PCIe PHY	单层，≤1024	Helios
UB 灵衢	华为	私有	多层，万卡级	CloudMatrix 384
HSL	海光	私有或 PCIe/以太网 PHY	多层	scaleX640
ALink System	阿里云/信通院	以太网 PHY	多层，≤2000	磐久 128
ETH-X	腾讯/信通院	以太网	单层，≤512	ETH64
EthLink	字节跳动	以太网	单层	大禹
高通量 ETH+	阿里云/中科院计算所	以太网	多层，≤1024	-

数据来源：《超节点技术体系白皮书》，鲜枣课堂，公司官网，东吴证券研究所

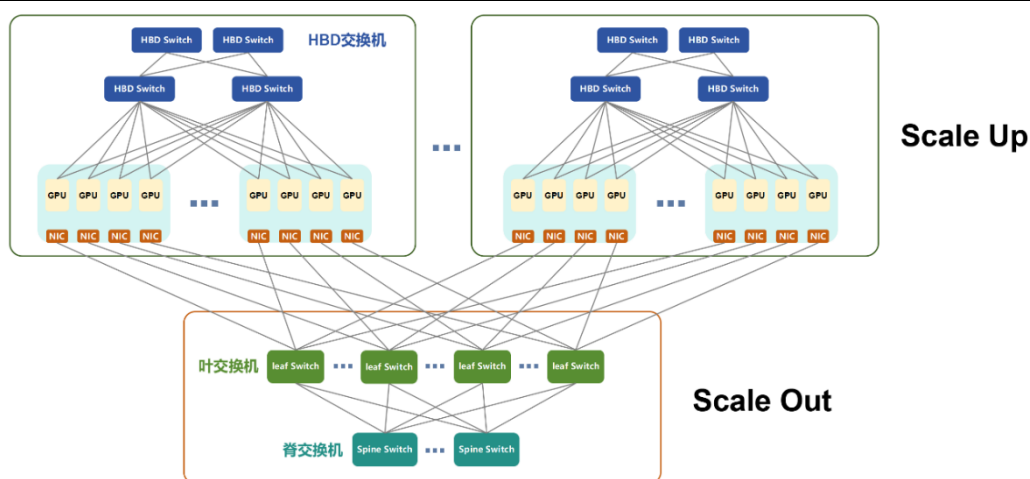
国产总线化路线更强调通过系统互联扩大有效算力供给。华为灵衢 UB、海光 HSL 等方案将统一编址、内存语义和多层组网纳入国产算力体系，核心目标是在单卡性能仍受约束的阶段，通过更大的协同计算域提升集群效率。该路线的产业意义在于把芯片、服务器、网络 and 系统软件连接成更完整的国产基础设施供给。

云厂商和开放互联路线则更关注负载适配和生态参与空间。阿里 ALink 从云侧训练、推理和资源池化需求出发，围绕 CPU/GPU/交换节点解耦重组系统架构；UALink、EthLink、ETH-X 等路线尝试复用以太网生态，为多芯片、多系统厂商提供共同接口。随着 Scale Up 从闭环产品走向多路线并行，交换芯片、连接器、线缆、光互联和整柜集成环节的参与空间有望扩大。

5.1.3. Scale Out: 超节点之间的横向扩展决定集群效率

Scale Out 决定超节点能否从单个高性能单元扩展为大规模算力基础设施。Scale Up 负责节点内部的高频通信，Scale Out 则连接多个高带宽域，承接数据并行、流水并行、存储访问、任务调度和故障恢复等需求。当单个高带宽域受到功耗、散热、连接复杂度和成本约束后，继续提升总算力需要依赖多个超节点横向扩展。

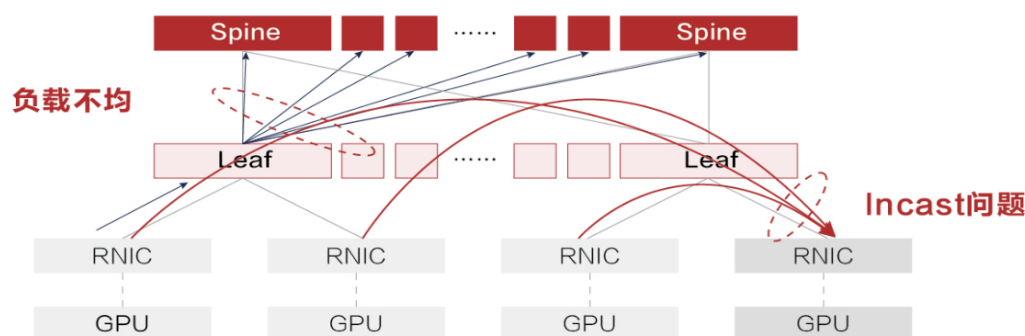
图8: Scale up & Scale out 组网示意图



数据来源：鲜枣课堂，东吴证券研究所

大规模集群的效率取决于网络在高负载下维持稳定吞吐的能力。训练任务中的同步通信具有突发性，网络拥塞、尾时延和丢包重传都会拖慢整体作业；推理侧多实例部署也会放大负载均衡和故障恢复压力。因此，Scale Out 能力不能只看端口速率，还要看 RDMA 效率、拥塞控制、路径调度和集群管理能力。

图9: AI 流量在集群中表现的负载不均和拥塞问题



数据来源：新华三官网，东吴证券研究所

Scale Out 的产业价值，在于把整柜超节点扩展为可持续运行的 AI 基础设施。网络设备、网卡、光模块、高速线缆和集群管理软件共同决定多超节点系统的可用性。随着产品形态从单柜走向 SuperPOD、SuperCluster 和 AI Factory，Scale Out 将从通用数据中心网络能力升级为超节点规模化部署的关键能力。

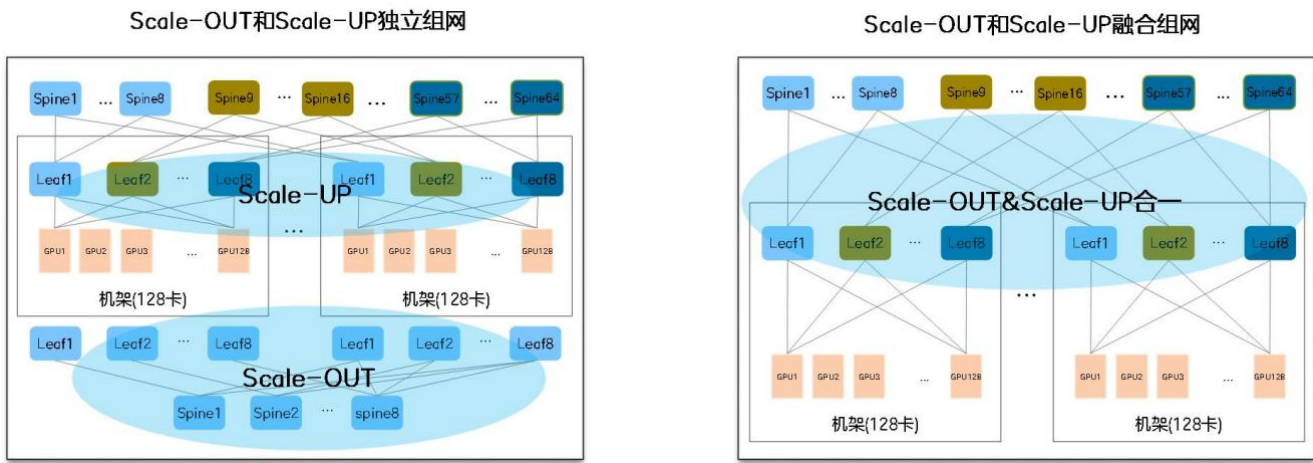
5.1.4. Scale Up 与 Scale Out 边界变化：从服务器边界转向性能边界

Scale Up 与 Scale Out 的边界正在从物理服务器边界转向性能边界。过去，Scale Up 主要发生在单机内部，Scale Out 主要发生在服务器之间；超节点把 Scale Up 扩展到整柜甚至多柜范围，使部分高频流量重新进入更高带宽、更低时延的协同域。架构划分的核心依据，逐步从“设备位于哪里”转向“流量需要怎样的通信能力”。

更大的 Scale Up 域可以提升模型并行和低频同步效率，也会带来更高连接密度、更复杂的交换结构和更严格的散热供电要求。高带宽域扩张不是单纯扩大卡数，而是通信、结构、供电、液冷和软件运维能力的同步升级。超节点的系统边界，最终取决于性能收益与工程复杂度的平衡。

分层互联将成为超节点演进的重要形态。柜内通过 Scale Up 承接最高频通信，柜间通过 Scale Out 承接横向扩展，POD 级系统中两类网络进一步融合。英伟达从 NVL72 走向 POD，华为从 SuperPoD 走向 SuperCluster，中科曙光从 scaleX640 走向万卡超集群，均体现出这一产品演进方向。

图10: Scale-Up 和 Scale-Out 融合和独立组网对比



数据来源:《中兴通讯超节点技术白皮书》, 东吴证券研究所

5.1.5. 光互联与 CPO: 高带宽域扩大后的下一阶段通信增量

光互联的增量来自高带宽域扩大后的连接压力。单柜内部短距互联中, 铜连接仍具备成本、功耗和可靠性优势, 尤其适合计算托盘与 Switch Tray 之间的高密度连接。当 Scale Up 向跨柜、多机架和 POD 级系统延伸后, 传输距离、线缆密度和维护难度上升, 光互联的重要性随之提升。

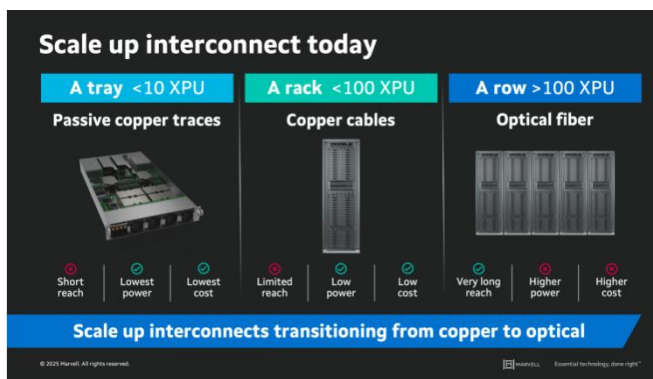
图11: 英伟达 GB200 内部使用铜连接



数据来源: 英伟达 GTC, 东吴证券研究所

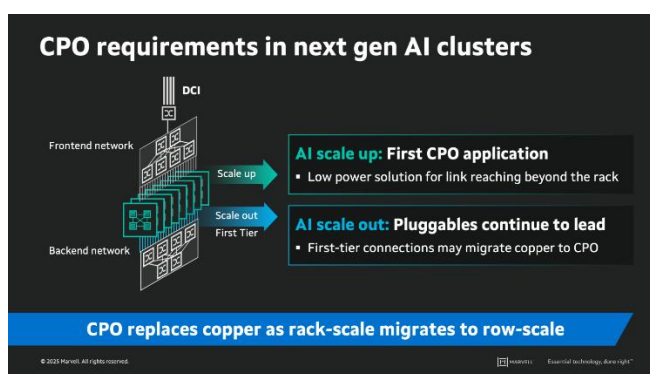
CPO 和硅光主要对应高容量交换系统的功耗与密度约束。随着交换芯片容量提升，传统可插拔光模块和高速电走线在面板密度、信号完整性和功耗上面临更大压力。CPO 将光电转换靠近交换芯片，缩短高速电连接距离，提升交换系统在大带宽场景下的能效和集成密度。

图12: Scale-up 互联由铜互联向光互联迁移



数据来源: Marvell 官网, 东吴证券研究所

图13: CPO 有望率先在跨机柜 Scale-up 场景落地



数据来源: Marvell 官网, 东吴证券研究所

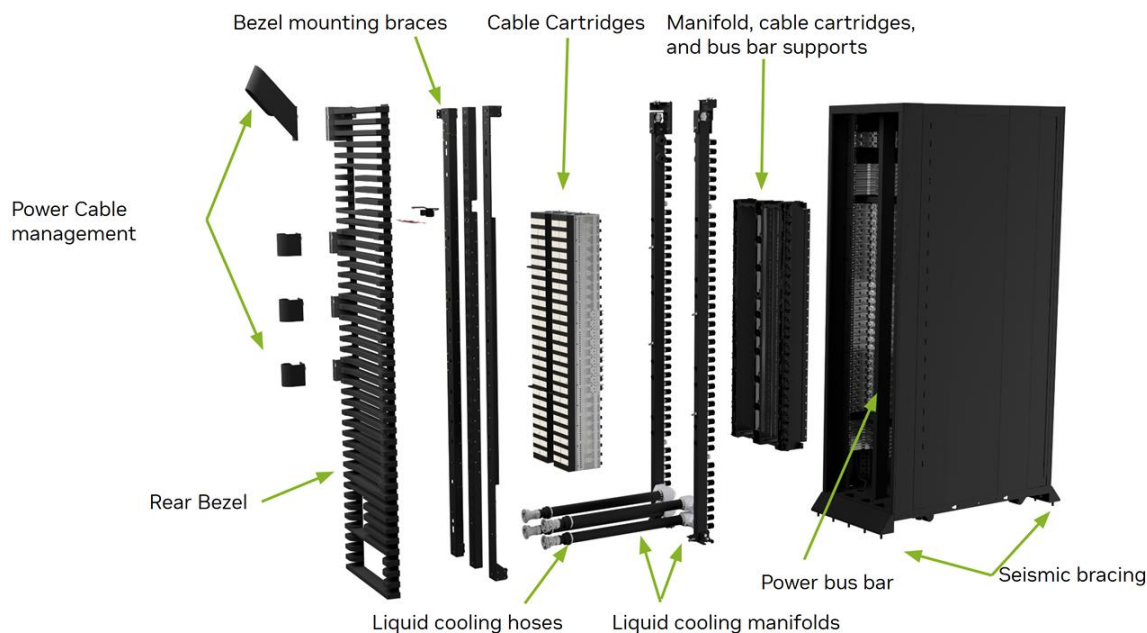
通信链路从柜内铜连接延伸至跨柜光互联，将带动光模块、硅光、CPO、连接器和系统集成环节价值提升。短期看，铜连接仍是整柜内部的主流选择；中长期看，随着超节点从单柜走向集群化部署，光互联将成为通信系统升级的重要增量。

5.2. 工程系统是超节点落地的第二主线：整柜级交付成为产品化门槛

工程系统决定超节点能否从高性能架构走向可复制交付。高带宽互联解决计算模组之间的协同效率，整柜工程则决定这一架构能否在高功率、高密度和高连接复杂度下长期稳定运行。超节点把计算、交换、供电、液冷和管理模块压缩到整柜乃至多柜空间内，产品难点从单点性能扩展到散热、电力、结构、可靠性和现场维护。

整柜级交付正在重塑服务器系统的价值构成。普通 AI 服务器阶段，厂商能力主要围绕单机主板、电源、散热和整机测试展开；超节点阶段，计算托盘、Switch Tray、Power Shelf、Busbar、液冷管路和系统软件需要作为统一系统协同设计。服务器厂商的能力边界由此从单机集成延伸到整柜系统工程，价值留存也从整机装配扩展到整柜集成、系统测试和现场交付。

图14: 整柜系统集成覆盖液冷、电力、线缆与结构工程

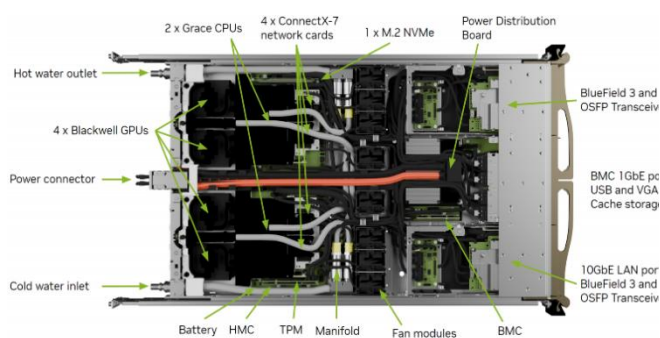


数据来源：英伟达官网，东吴证券研究所

5.2.1. 散热：从服务器级散热走向整柜级液冷

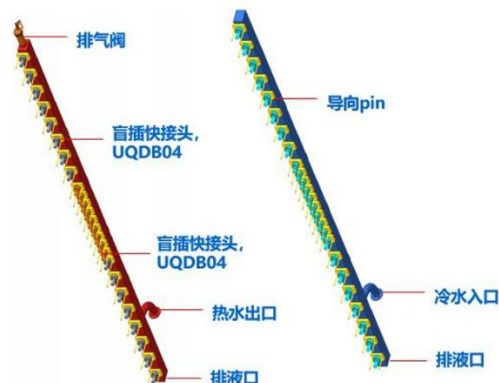
散热能力正在成为超节点规模化部署的硬约束。超节点把计算节点、交换节点和高速连接集中到整柜内，热源密度更高，负载持续时间更长，传统风冷在效率、噪声和空间占用上的边界更快显现。随着单柜功率上行，液冷从高端配置转向高功率密度系统的基础设施，冷板、CDU、Manifold、快接头和漏液检测共同构成整柜热管理体系。整柜液冷的价值不只是降低芯片温度，而是支撑高密系统稳定交付，核心难点在于流量分配、节点插拔、漏液监控和现场维护。液冷系统由此从服务器散热部件升级为整柜工程能力，直接影响超节点的可部署性和长期运行稳定性。

图15：GB200 液冷计算托盘内部结构



数据来源：英伟达官网，东吴证券研究所

图16：机柜背板通过 manifold 供水

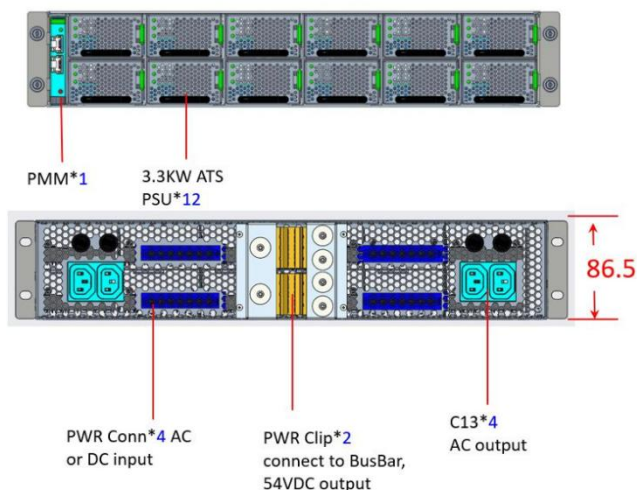


数据来源：ODCC《ETH-X 超节点 AI 整机柜设计规范》，东吴证券研究所

5.2.2. 供电：从电源模块走向整柜级高功率供电

供电系统正在从服务器内部电源配置升级为整柜级功率平台。超节点需要同时支撑计算节点、交换节点、管理系统和液冷系统，整柜功耗更高，负载波动更集中，Power Shelf、Busbar、BBU 和冗余架构不再只是配套部件，而是决定系统能否稳定满载运行的基础环节。单柜功率继续上行后，电源转换效率、母线连接可靠性、功耗监控和局部故障保护都会影响整柜可用性。服务器厂商若要承接超节点交付，需要同时理解机柜内部供电设计和客户机房侧供配电条件，供电能力也将成为整柜项目落地的重要约束。

图17: 机柜功耗持续攀升带动 Power shelf 和 PSU 需求



数据来源: ODCC《ETH-X 超节点 AI 整机柜设计规范》, 东吴证券研究所

图18: 采用 Busbar 为各节点提供电源

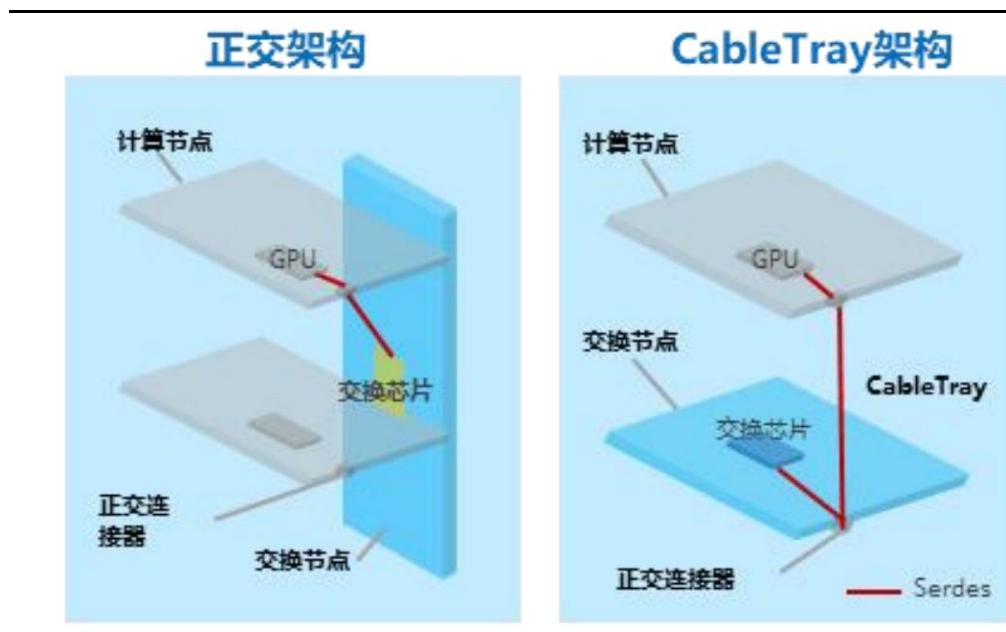


数据来源: ODCC《ETH-X 超节点 AI 整机柜设计规范》, 东吴证券研究所

5.2.3. 结构与线缆: 高密度系统中的机械工程问题

结构与线缆设计决定高密度超节点能否实现可制造、可部署和可维护。超节点机柜内部同时容纳计算托盘、Switch Tray、供电通道、液冷管路和高速连接，空间布局稍有不合理，就会带来线缆过长、连接器应力增加、散热路径受阻或节点维护困难。Scale Up 域扩大后，计算节点与交换节点之间需要依赖大量高速连接保持低损耗传输，Cable Tray、背板/中板、正交连接和托盘化布局的意义在于压缩连接路径并提升装配一致性。线缆长度、弯折半径、连接器可靠性和 EMI/EMC 设计都会影响链路质量，结构设计由此从机械外壳问题上升为整柜高带宽能力释放的关键环节。

图19: OEX 与 Cable Tray 方案对比

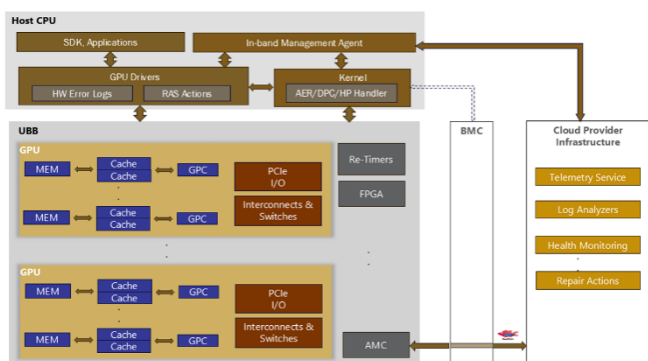


数据来源：《中兴通讯超节点技术白皮书》，东吴证券研究所

5.2.4. 可靠性：从单机可用性走向系统级容错

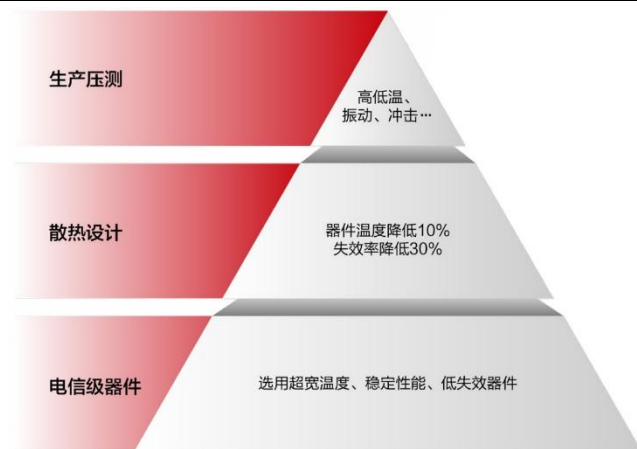
超节点可靠性问题来自系统规模扩大后的故障放大效应。单台服务器故障通常影响一个节点，超节点内部计算、交换、供电、冷却和管理系统高度耦合，链路异常、温度异常或供电异常都可能影响正在运行的训练和推理任务。节点数量和连接数量增加后，故障不再是偶发事件，而是系统设计必须持续吸收的常态。可靠性能力需要从单机测试升级为整柜级 RAS 体系，链路重传、故障隔离、健康监控、热插拔和自动恢复决定系统韧性；整柜烧机、压力测试、温度循环、液冷可靠性和供电冗余测试决定产品能否从工程样机走向商业部署。

图20：超节点系统级 RAS 管理架构



数据来源：OCP，东吴证券研究所

图21：硬件可靠性策略

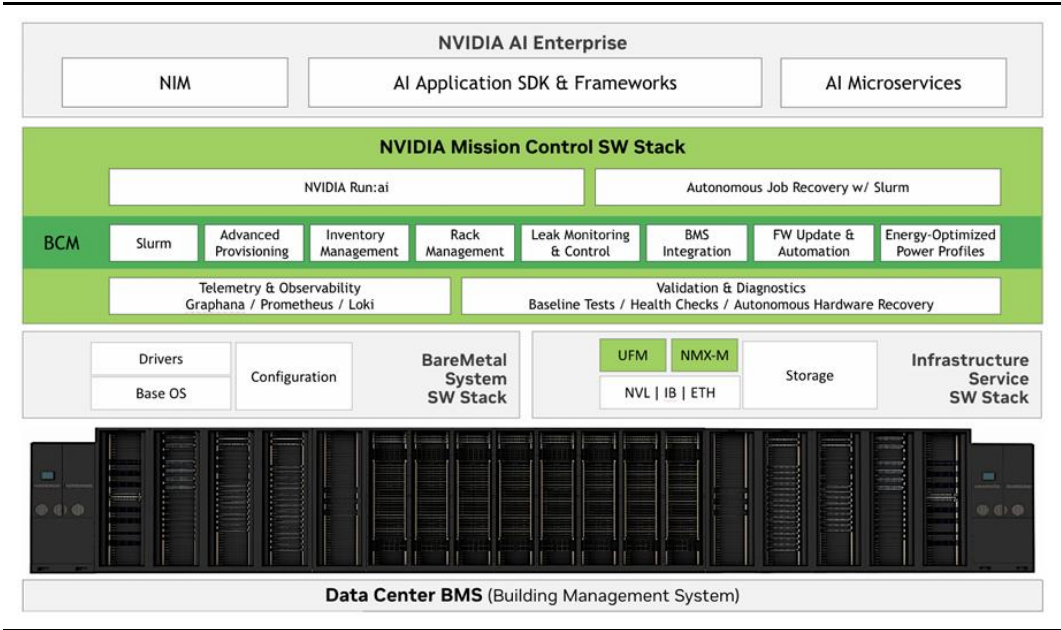


数据来源：《超节点发展报告》，东吴证券研究所

5.2.5. 软件与运维：硬件互联决定上限，系统软件决定释放效率

硬件互联提供性能上限，系统软件决定上限能够释放多少。超节点内部节点多、链路多、状态多，单靠硬件连接无法自动形成稳定算力，通信库、驱动/固件、资源调度、作业编排和集群监控，需要把计算、交换、供电和散热系统组织成可持续运行的基础设施单元。训练任务需要围绕模型并行和通信拓扑进行资源匹配，推理服务则需要围绕上下文长度、KV Cache、并发请求和服务时延进行动态调度。随着超节点从整柜走向 POD 和集群，拓扑发现、固件管理、性能监控、故障预测和作业恢复将成为规模化交付的必要能力。

图22：系统软件支撑集群调度、监控与恢复



数据来源：英伟达官网，东吴证券研究所

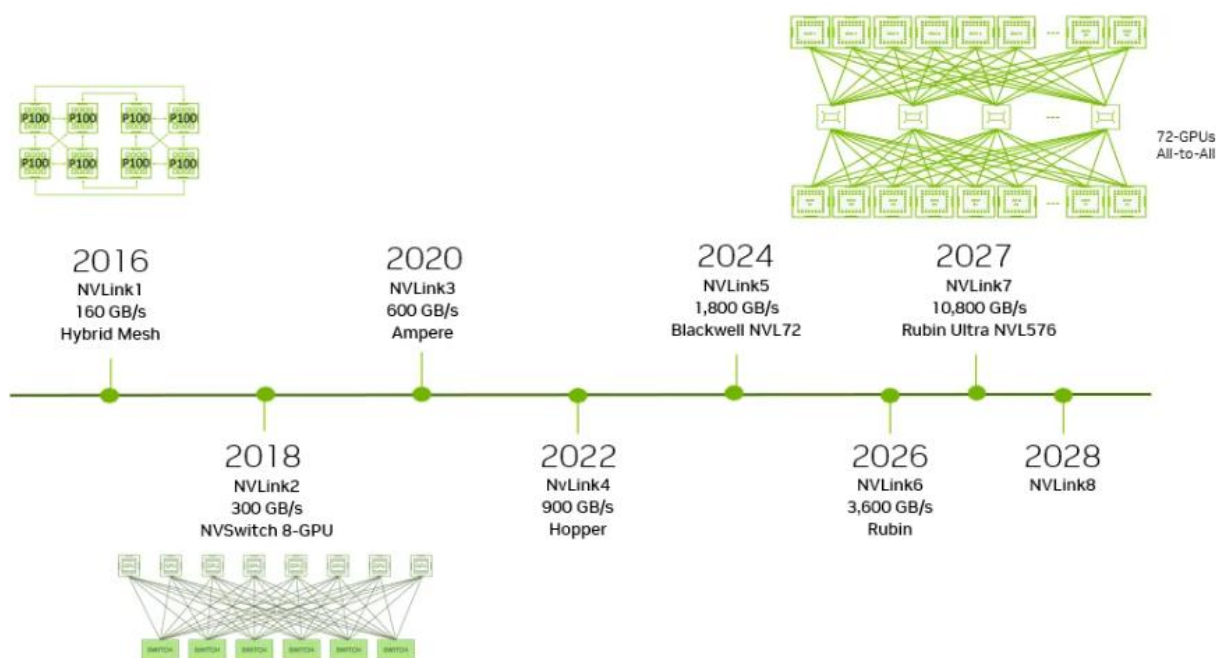
6. 主流超节点产品演进：闭环生态、国产体系与开放路线并行

6.1. 英伟达：从 NVLink 整柜系统走向 AI Factory

英伟达超节点路线的核心，是以 NVLink/NVSwitch 为主轴，将 GPU 服务器升级为整柜级 Scale Up 系统。NVLink 自 2016 年启动演进，2018 年引入 NVSwitch 后，互联能力从 GPU 点对点连接逐步走向交换式高带宽域；Blackwell NVL72、Rubin NVL72 和 Rubin Ultra NVL576 则进一步把这一架构固化为标准化整柜产品。相比单一 GPU 代际升级，英伟达更强的地方在于 GPU、互联协议、交换芯片和系统软件可以同步迭代，形成清晰的产品节奏。

图23：NVLink 的 Scale up 域迅速扩张

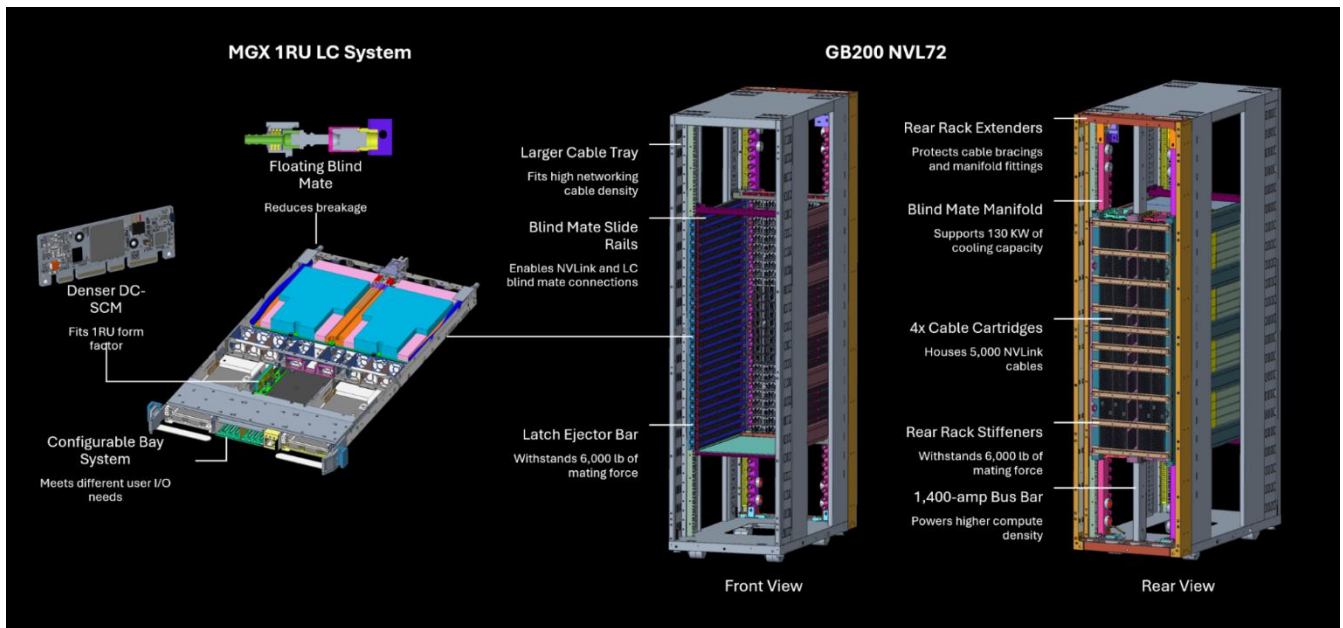
NVLink - Scale-Up Maturity with Rapid Innovation Pace



数据来源：英伟达官网，东吴证券研究所

从 NVL72 开始，英伟达的产品边界已经从整柜系统继续推向 AI Factory。单柜 NVLink 域负责承接高频通信，InfiniBand 和 Spectrum-X 负责多柜扩展，Mission Control 等工具则承担部署、诊断、故障恢复和工作负载管理。英伟达路线的关键变化在于，客户采购对象不再只是 GPU 服务器，而是可以持续扩展和统一运维的 AI 基础设施单元。

图24: GB200 NVL72 由计算托盘走向整柜级 Scale Up 系统

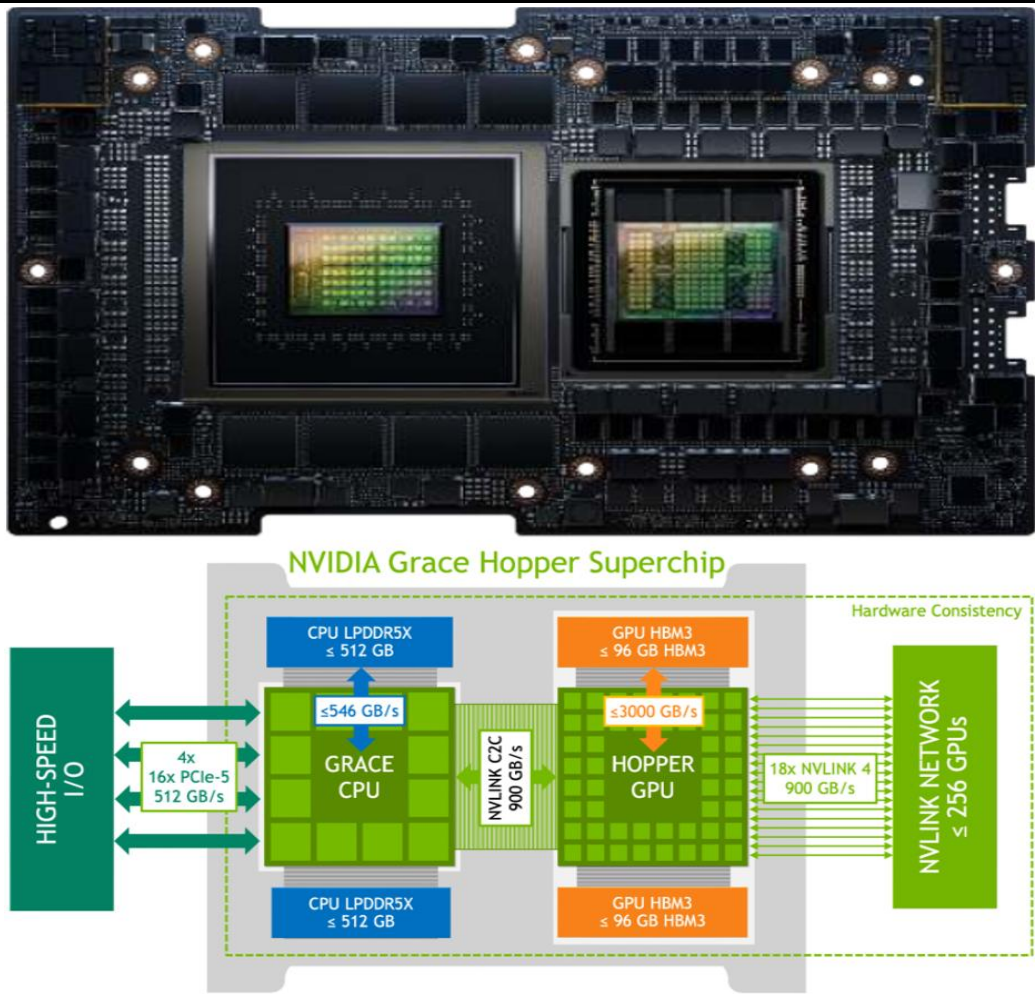


数据来源：英伟达官网，东吴证券研究所

6.1.1. GB200 NVL72 与 GB300 NVL72: 整柜级超节点形态确立并向推理增强

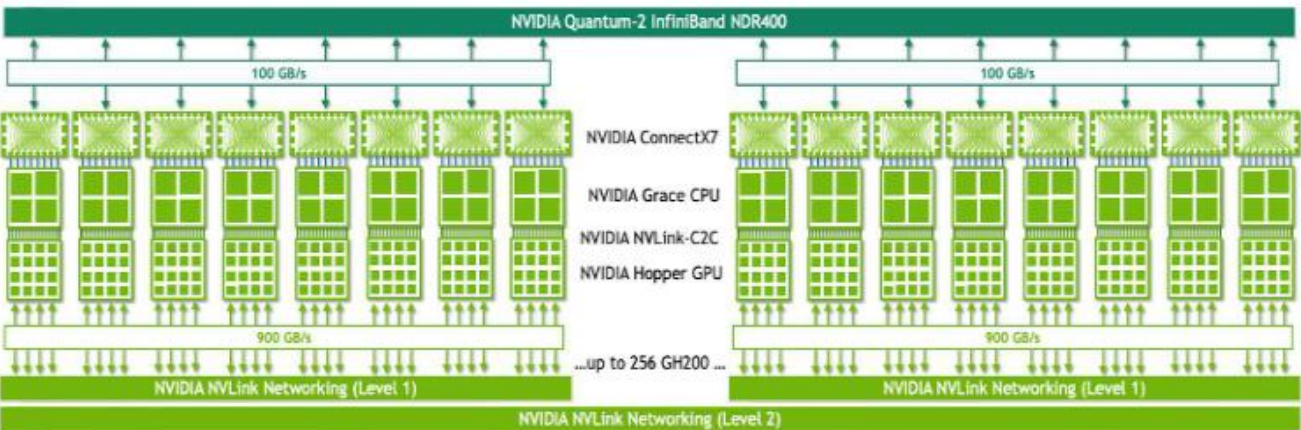
GH200 阶段为英伟达超节点产品化提供了早期雏形。DGX GH200 和 GH200 NVL32 已经将 Grace CPU、Hopper GPU、NVLink-C2C 和 NVLink Switch 放入同一系统，使 GPU 服务器开始具备跨节点协同和大共享内存特征。该阶段的意义在于确立了从节点级产品走向系统级产品的方向，但其形态仍更接近 AI supercomputer 或系统参考架构，尚未形成 NVL72 这种标准化整柜单元。

图25：英伟达首款 CPU-GPU 超级芯片



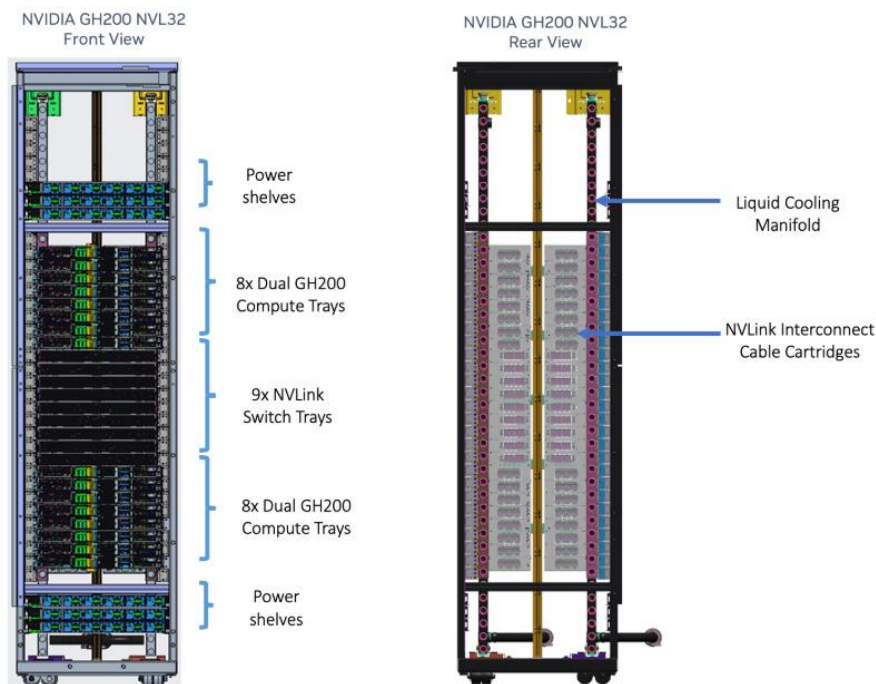
数据来源：英伟达官网，东吴证券研究所

图26：DGX GH200 超级计算机通过 NVLink 连接最多 256 个 GH200 超级芯片



数据来源：英伟达官网，东吴证券研究所

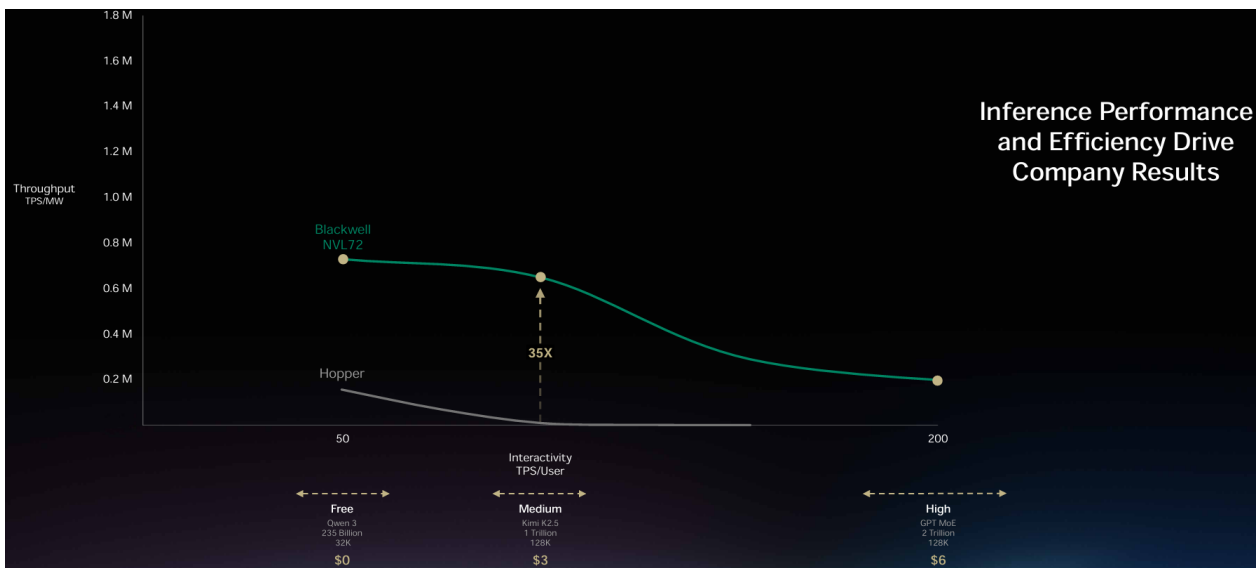
图27: GH200 NVL32 单柜级 NVLink 系统结构



数据来源：英伟达官网，东吴证券研究所

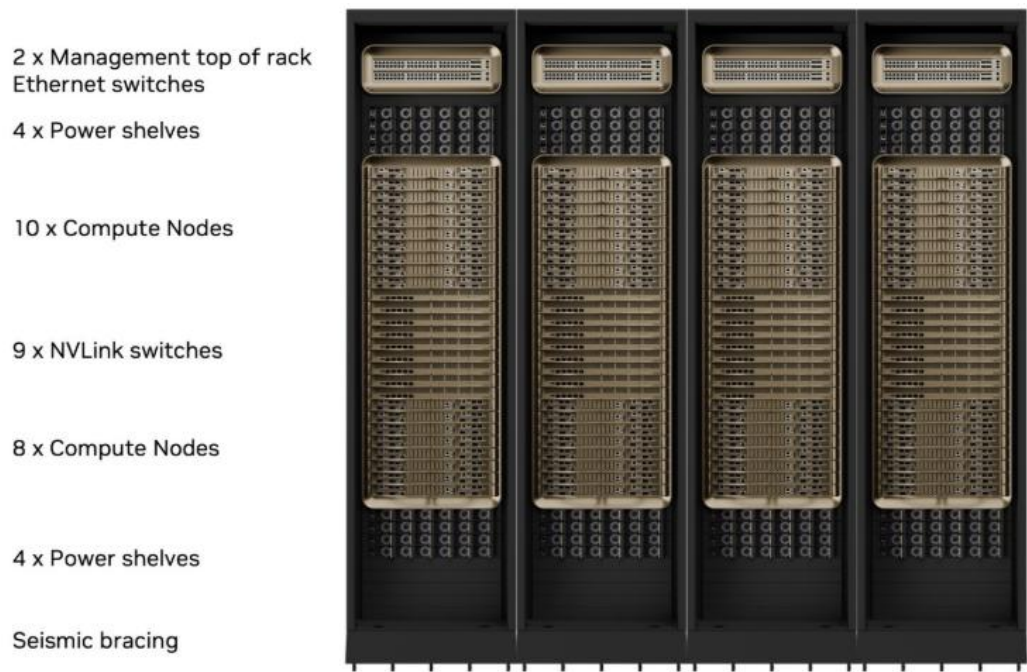
GB200 NVL72 标志着英伟达将超节点稳定到整柜级 72 GPU 高带宽域，并实现性能口径和工程形态同步成型。该系统把 36 颗 Grace CPU 与 72 颗 Blackwell GPU 放入同一机柜，通过 18 个 Compute Tray、9 个 Switch Tray、整柜液冷和柜内高速连接形成 rack-scale 系统；单柜 13.4TB HBM3E、576TB/s HBM 带宽和 130TB/s NVLink 带宽，为训练和推理提供了明确的系统能力边界。相比 GH200，GB200 的关键变化在于产品边界从服务器扩展到机柜，客户部署基础单元从单机设备升级为整柜计算域，复杂互联结构也被固化为可复制、可测试、可交付的标准产品。

图28: Blackwell NVL72 推理吞吐与能效较 Hopper 显著提升



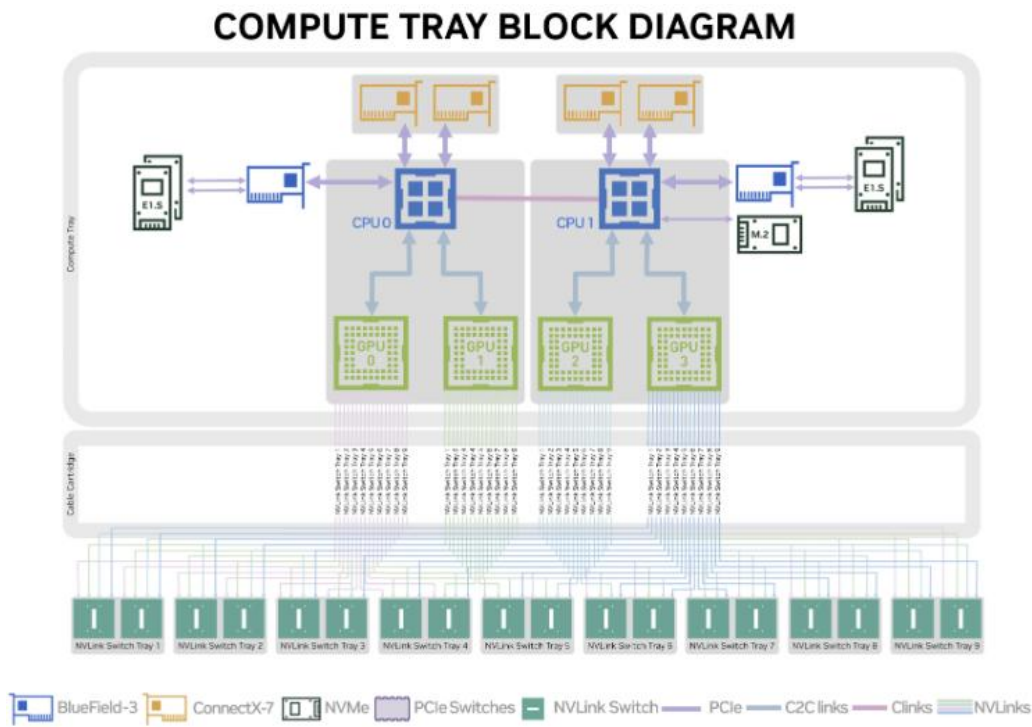
数据来源：英伟达官网，东吴证券研究所

图29：GB200 NVL72 以整柜级模块化布局确立 72 GPU 标准交付形态



数据来源：英伟达官网，东吴证券研究所

图30：GB200 NVL72 以 2 CPU + 4 GPU 计算托盘为基本单元



数据来源：英伟达官网，东吴证券研究所

GB300 NVL72 延续 GB200 的整柜架构，产品重心进一步转向推理、长上下文和 AI reasoning。GB300 仍采用 72 GPU 加 36 Grace CPU 的 rack-scale 形态，GPU 平台升级为 Blackwell Ultra，并强化 HBM 容量、注意力计算和推理吞吐。37TB fast memory、20TB GPU memory 和 130TB/s NVLink 带宽等产品口径，指向更高并发、更长上下文和更高单位机柜 token 产出。GB200 确立整柜级超节点形态，GB300 则在同一形态下提高推理生产效率。

图31: GB300 NVL72 网络带宽较 GB200 NVL72 提升 2 倍，FP4 推理性能提升 1.5 倍

	GB300 NVL72	vs. GB200 NVL72	vs. HGX H100
FP4 Inference ¹	1.4 1.1 ExaFLOPS	1.5x	70x
HBM Memory	20 TB	1.5x	30x
Fast Memory	40 TB	1.3x	65x
Networking Bandwidth	14.4 TB/s	2x	20x

Table 1. NVIDIA Blackwell Ultra specifications compared to NVIDIA GB200 NVL72 and NVIDIA HGX H100

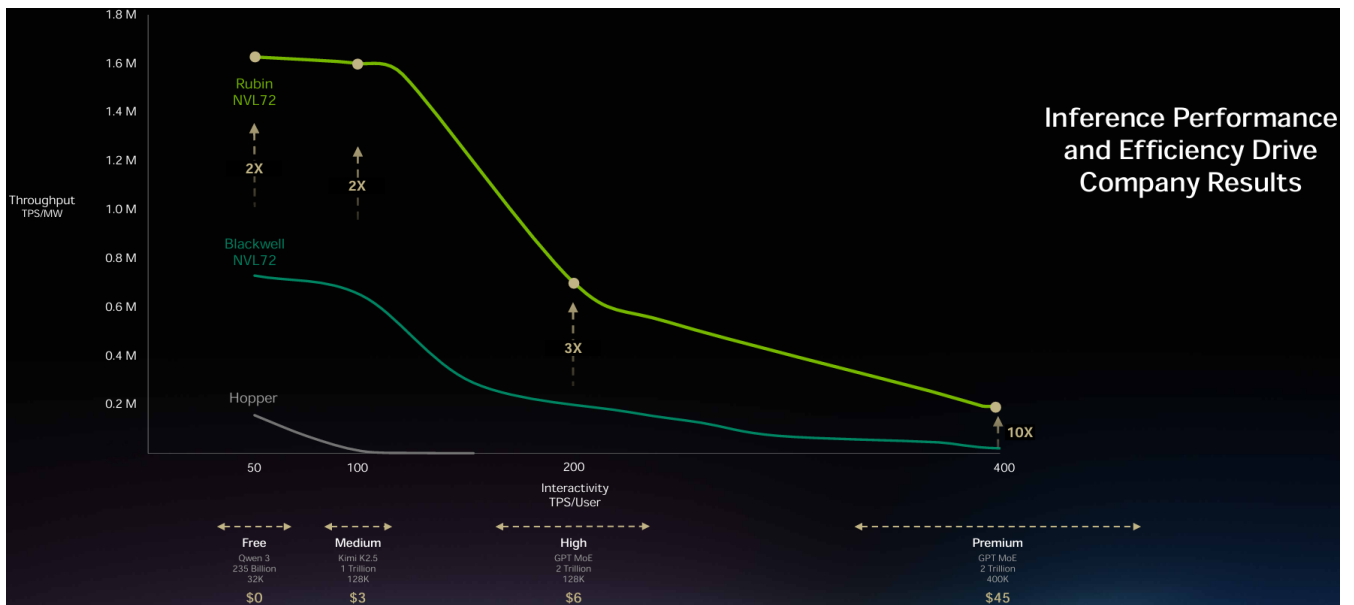
¹With Sparsity | Without Sparsity

数据来源：英伟达官网，东吴证券研究所

6.1.2. Vera Rubin、Rubin Ultra 与 Kyber Rack: 从单柜 NVLink 域走向 POD 级 AI Factory

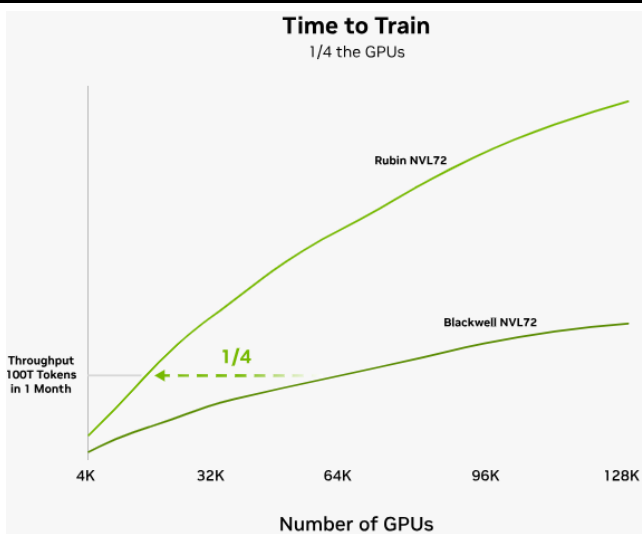
Vera Rubin NVL72 延续 NVL72 的单柜产品形态，代际重心转向更高的 AI Factory 系统效率和更强的工程可交付性。该系统仍以 72 个 Rubin GPU 和 36 个 Vera CPU 组成单柜计算单元，并升级至 HBM4 内存体系；官方产品口径显示，Vera Rubin NVL72 对应 20.7TB HBM4、260TB/s NVLink 带宽，Scale Up 带宽较 Blackwell NVL72 翻倍。第三代 MGX 机架、模块化托盘和更少线缆设计，指向更快部署、更低维护复杂度和更高现场交付一致性；NVLink 交换托盘支持维护模式和在线更换，也说明整柜产品竞争已经从“连接更多 GPU”推进到“稳定交付高密度系统”。

图32: Rubin NVL72 推理性能与能效进一步提升



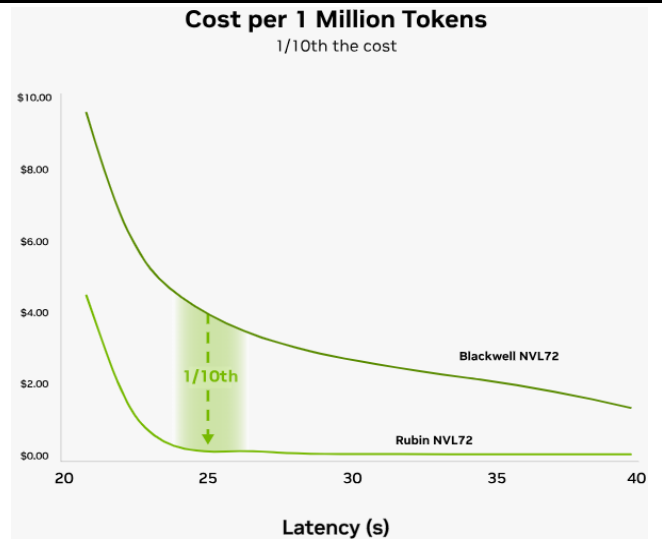
数据来源: 英伟达官网, 东吴证券研究所

图33: Rubin NVL72 以更少 GPU 完成同等训练任务



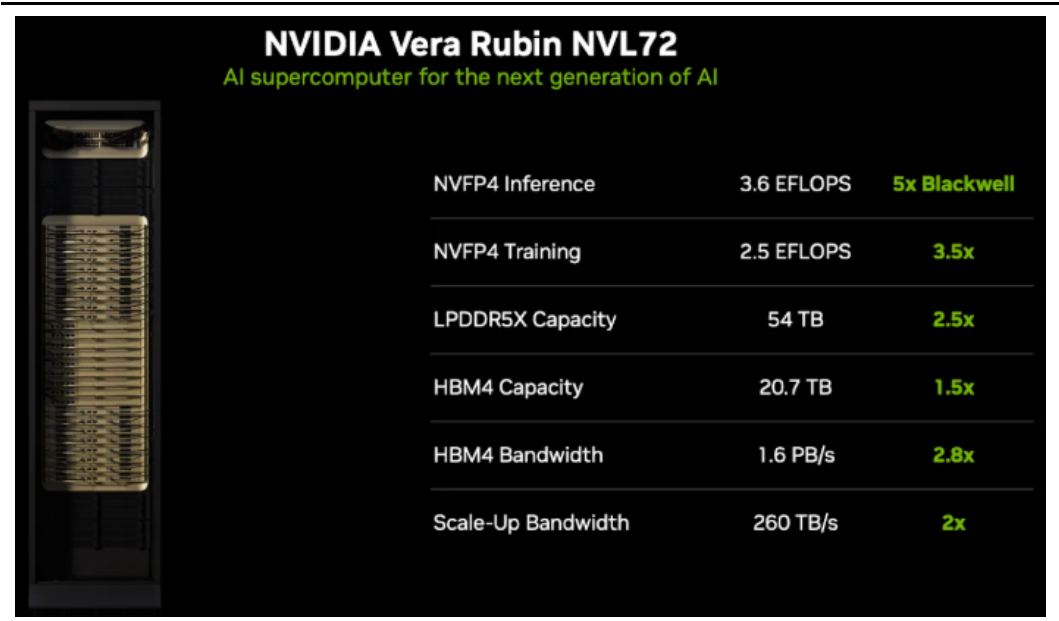
数据来源: 英伟达官网, 东吴证券研究所

图34: Rubin 推理成本降至 Blackwell 的十分之一



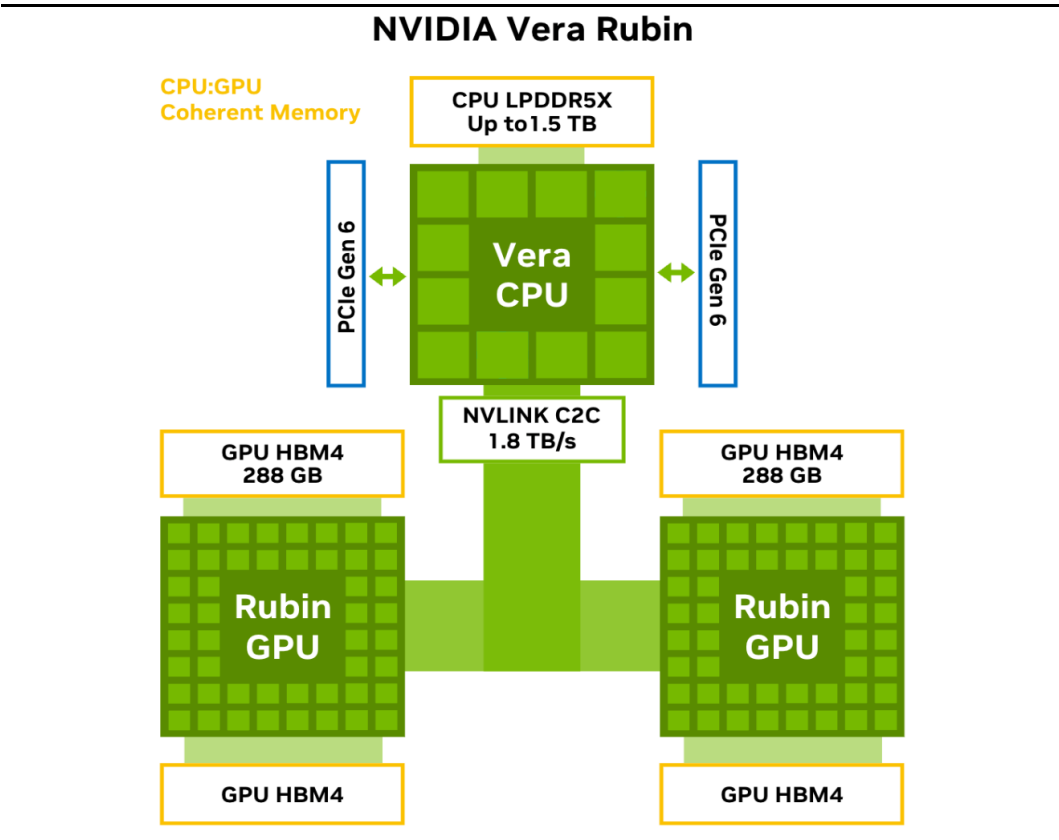
数据来源: 英伟达, 东吴证券研究所

图35: Vera Rubin NVFP4 推理性能达 3.6EFLOPS，较 Blackwell 提升 5 倍



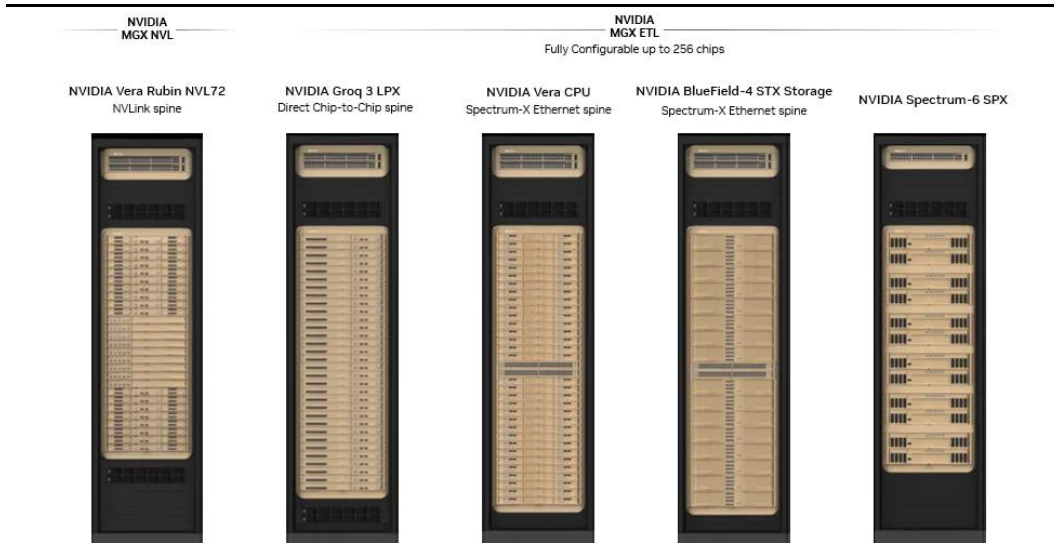
数据来源：英伟达官网，东吴证券研究所

图36: Vera Rubin 进一步强化 CPU-GPU 一致性内存与 NVLink-C2C 互联



数据来源：英伟达官网，东吴证券研究所

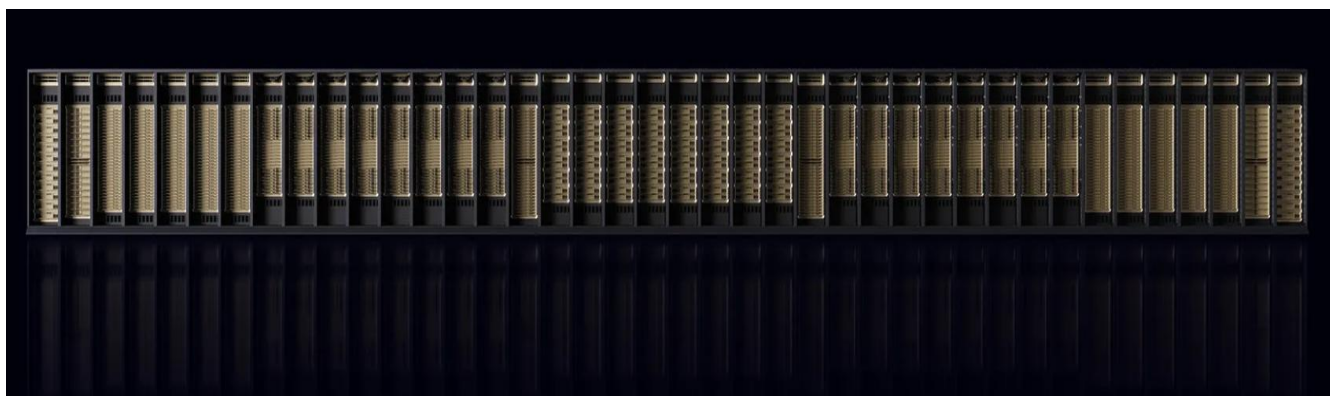
图37: POD 以五类机架系统将 NVL72 扩展为多机架 AI Factory 单元



数据来源: 英伟达官网, 东吴证券研究所

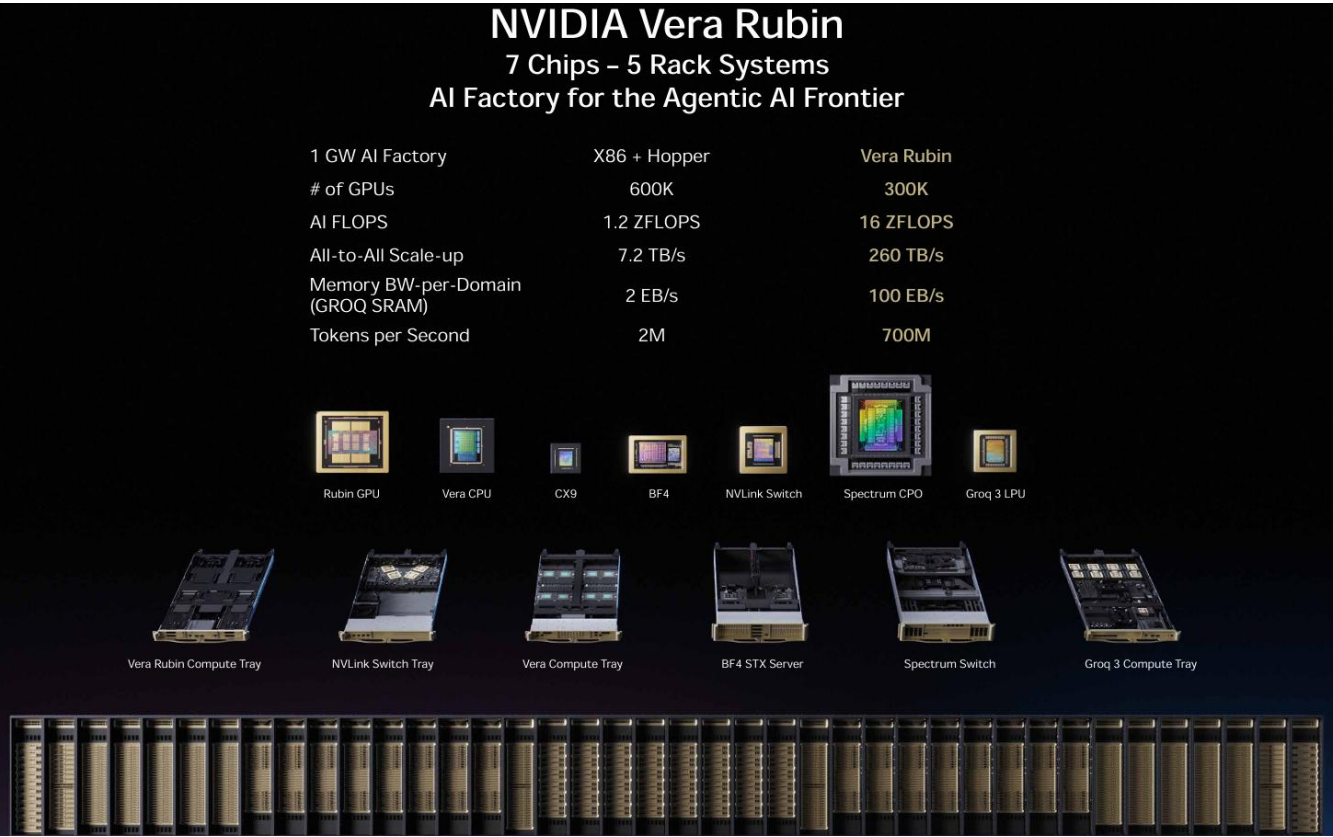
Vera Rubin POD 将 NVL72 从单柜产品扩展为多机架 AI Factory 单元。POD 级形态由 GPU 计算、CPU、网络、存储和低时延推理等多类机架组成, 产品边界从标准化整柜延伸到围绕 AI Factory 工作负载组织的系统组合。Vera Rubin POD 由 40 个机架、1,152 个 Rubin GPU 组成, 并形成 60EFLOPS 和 10PB/s 总 Scale Up 带宽的系统口径, 说明英伟达正在把 NVL72 从单柜基础单元推进为可组合的多机架算力工厂。

图38: Vera Rubin POD 由 40 个机架、1,152 个 Rubin GPU 组成



数据来源: 英伟达官网, 东吴证券研究所

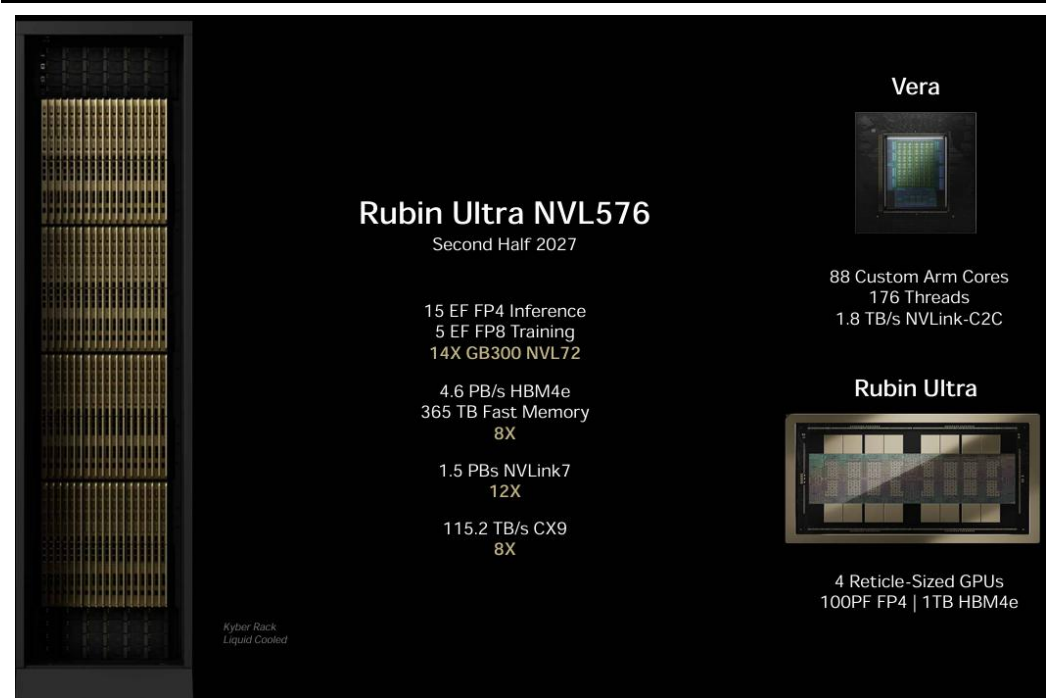
图39: Vera Rubin 通过“七芯片、五机架”协同，将 AI Factory 竞争推进到系统级 token 产出效率



数据来源：英伟达官网，东吴证券研究所

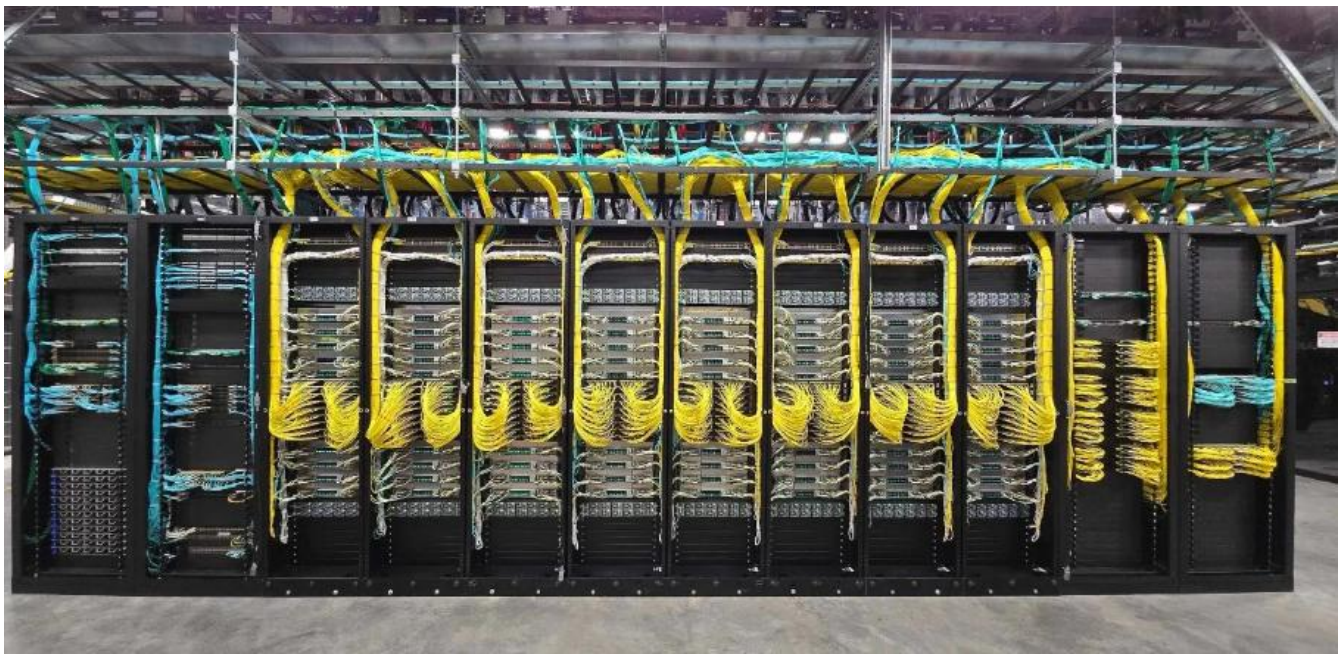
Rubin Ultra 与 Kyber Rack 代表英伟达继续扩大 NVLink 域的下一阶段方向。 Rubin Ultra NVL576 将多个 MGX NVL 机架组合进更大的 Scale Up 系统，Kyber Rack 则进一步指向 NVL1152 等更大规模架构。随着高带宽域从单柜扩展到多机架，产品竞争也从单柜性能推进到跨机架互联、系统管理和数据中心设施匹配。英伟达路线由此从芯片平台竞争继续延伸到 AI Factory 级系统竞争。

图40: Rubin Ultra NVL576 系统性能约为 GB300 NVL72 的 14 倍



数据来源：英伟达官网，东吴证券研究所

图41: 英伟达 Polyphe 原型机，基于 GB200 的多机架 NVL576 架构



数据来源：英伟达官网，东吴证券研究所

图42: Kyber 是下一代 MGX 机架，将首先应用于 Rubin Ultra



数据来源：英伟达官网，东吴证券研究所

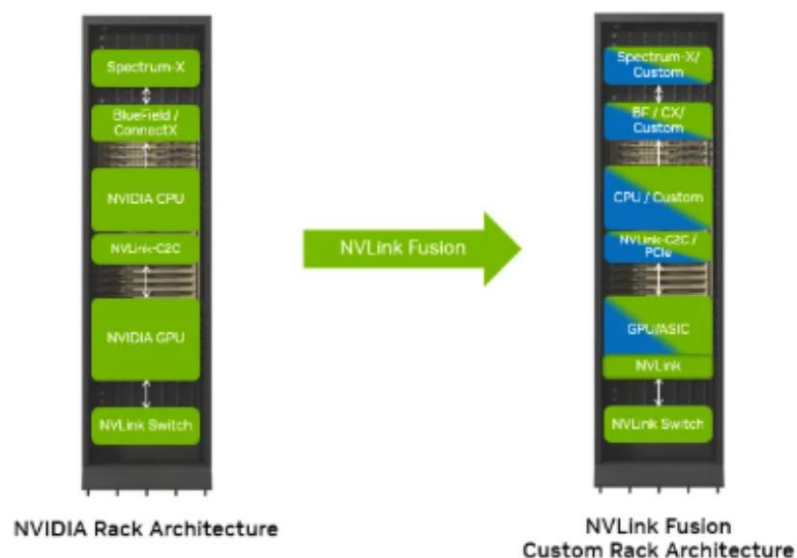
6.2. 海外开放路线：AMD Helios 与 UALink 生态形成补充

海外开放路线的核心，是在 NVLink 闭环体系之外提供可复用的 rack-scale AI 系统路径。英伟达已经以 NVLink/NVSwitch、整柜液冷和软件生态形成强控制力，AMD Helios 与 UALink 则代表另一类产业方向：通过开放机架、开放 Scale Up 标准和以太网 Scale Out，降低非英伟达生态构建整柜级 AI 基础设施的门槛。其短期意义不在于直接替代英伟达产品节奏，而在于为非闭环生态提供共同接口和可复制整柜形态。

图43: NVLink Fusion 显示英伟达正将 NVLink 闭环能力向开放芯片生态延伸

NVLink Fusion

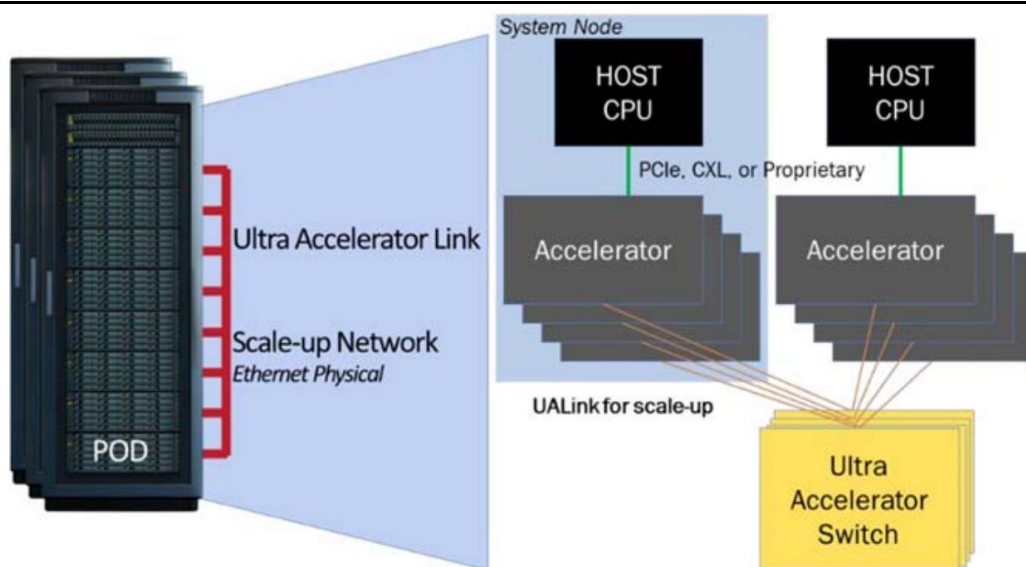
Custom Combinations with Open Standards Integration



数据来源：英伟达官网，东吴证券研究所

开放路线的产业价值集中在供应链参与空间和客户选择弹性。当 AI 基础设施从服务器走向整柜，客户关注的对象从单卡性能扩展到加速器互联、整柜液冷、网络扩展、管理软件和系统验证。开放机架和开放互联标准有助于降低生态绑定，使芯片、交换、线缆、光电器件和整柜系统围绕统一接口协同演进。该路线仍处于产品化推进阶段，后续关键在于 Helios 等整柜系统能否完成真实负载验证，并推动 UALink 从标准接口走向规模部署。

图44: UALink 面向 Scale Up 机架基础设施构建多节点、多平面交换网络



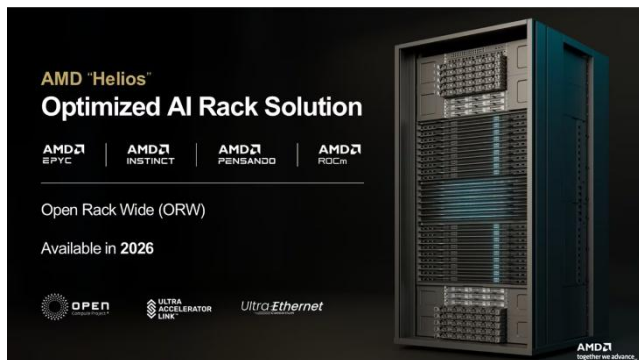
数据来源: UALink《UALink 200G 1.0 Scale-Up 互联技术白皮书》，东吴证券研究所

6.2.1. AMD Helios: 海外非英伟达生态进入 rack-scale 系统竞争

AMD Helios 的意义在于将海外非英伟达生态推进到整柜级系统竞争。过去非英伟达路线更多围绕单卡、服务器和局部平台能力展开，Helios 则以 MI450 系列 GPU、开放机架、液冷和以太网 Scale Out 构建 rack-scale 参考系统，产品方向对齐整柜级 AI 基础设施。其竞争对象不再只是单个加速器，而是由计算芯片、Scale Up 互联、Scale Out 网络和整柜工程共同构成的系统能力。

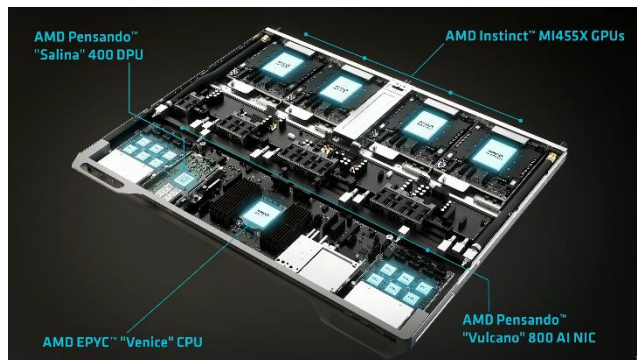
Helios 的产品定位体现出开放整柜路线对云厂商和系统厂商的吸引力。单柜 72 颗 MI450 系列 GPU、HBM4 内存体系、Scale Up 互联和以太网 Scale Out，使其具备与主流 rack-scale 系统同维度比较的产品基础。相比闭环路线，Helios 更强调开放机架和 OEM/ODM 伙伴可复用的参考设计，有利于客户根据负载类型、数据中心条件和供应链需求进行定制，也使液冷、高速连接、网络适配和系统验证环节获得更明确的参与空间。

图45: AMD Helios 以开放机架切入超节点系统竞争



数据来源: AMD 官网, 东吴证券研究所

图46: Helios 计算托盘将 CPU、GPU、DPU 与 AI NIC 集成在同一模块



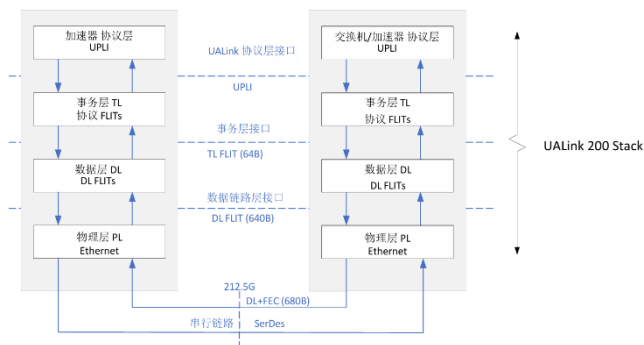
数据来源: CES 2026, 东吴证券研究所

Helios 的短期角色仍是海外开放路线的样板和补充。整柜级 AI 系统的竞争门槛不只在硬件规格,也在开发者生态、通信库、模型框架适配和客户部署经验。英伟达已经通过多代 NVLink 整柜产品形成较强产品节奏, Helios 仍需要经历从参考设计到批量交付、从硬件规格到软件生态、从实验室验证到真实负载运行的完整过程。其后续位置,将取决于 MI450 平台导入、UALink 生态成熟和云厂商规模部署节奏。

6.2.2. UALink: 开放 Scale Up 标准为非闭环生态提供接口

UALink 的核心价值,是为非 NVLink 生态提供统一的加速器 Scale Up 接口。超节点进入整柜级阶段后,加速器之间能否形成低时延、高带宽、可扩展的协同域,直接影响整柜系统的有效算力释放。UALink 200G 1.0 试图在开放生态中提供标准化互联基础,支持跨厂商加速器协同和多节点 Scale Up AIPOD 组网。对于 AMD、云厂商、芯片厂商和系统厂商而言,统一接口有助于减少互联碎片化,使不同 GPU/XPU、交换芯片、线缆和管理软件围绕共同标准形成系统组合。

图47: UALink 分层架构



数据来源: UALink 《UALink 200G 1.0 Scale-Up 互联技术白皮书》, 东吴证券研究所

图48: UALink 联盟董事会成员



数据来源: UALink 官网, 东吴证券研究所

UALink 推动开放路线从单点硬件竞争走向系统标准竞争。在专有互联体系中，GPU、交换、通信软件和整柜产品由单一主导方协同定义；开放体系则需要通过标准接口协调多方部件和软件生态。UALink 的产业意义在于把 Scale Up 能力从单一闭环系统中抽离出来，使整柜系统可以围绕开放互联构建更大的供应链协作空间。随着加速器规模从单柜走向多柜，交换芯片、端侧接口、背板/中板、高速线缆、光互联和整柜测试的重要性也会同步提升。

图49：UALink 演进时间线



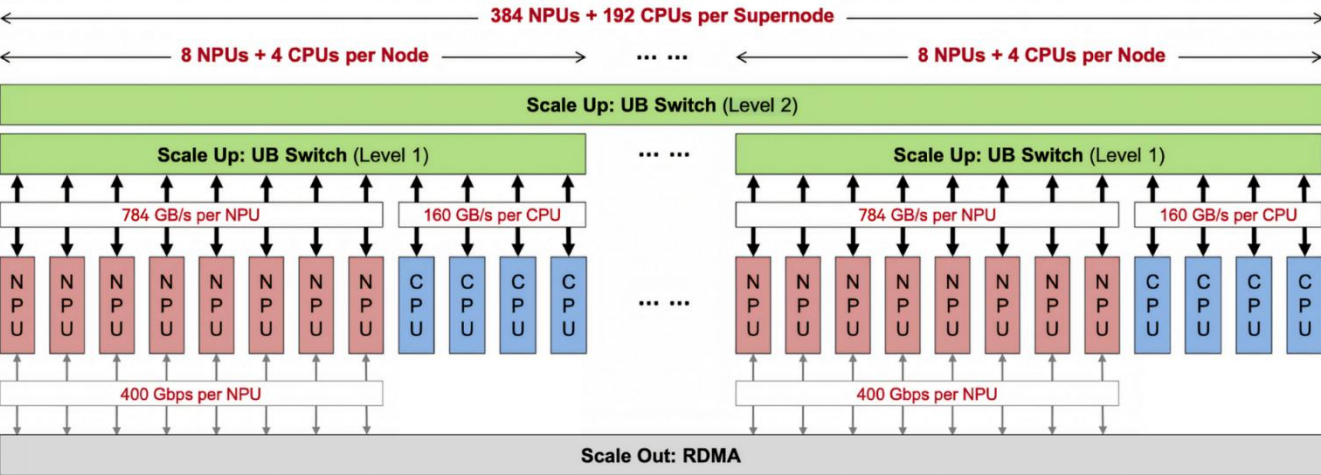
数据来源：2026 Open AI Infra Summit，东吴证券研究所

UALink 的约束在于标准落地需要芯片、软件和整柜验证共同成熟。Scale Up 对时延、顺序、链路可靠性和内存语义的要求高于传统数据中心网络，开放标准发布只能解决共同接口问题，产品竞争还需要经过芯片实现、互操作测试、系统管理适配和真实负载运行。海外开放路线已经形成清晰方向，但在产品成熟度、软件生态和交付确定性上仍需要追赶英伟达闭环体系。

6.3. 华为：以系统级扩展构建国产超节点体系

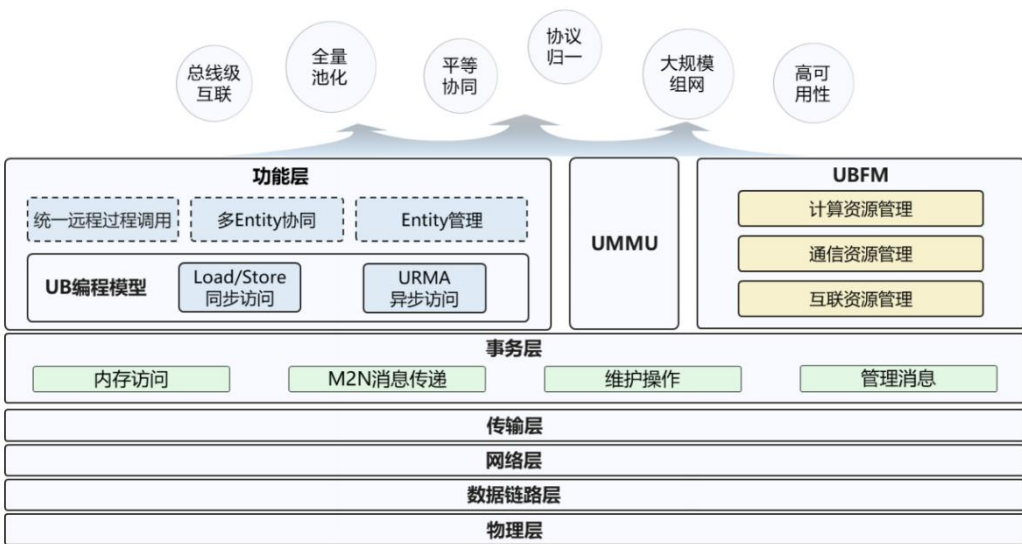
华为超节点路线的核心，是在国产算力供给约束下，通过系统级扩展提升有效算力。与英伟达以 NVL72 作为单柜 rack-scale 锚点不同，华为更强调以灵衢 UnifiedBus 组织更大的协同计算域，并通过 SuperPoD/SuperCluster 逐级放大系统规模。对等计算、统一编址和跨节点访存构成其互联基础，产品竞争重点落在如何把更多昇腾 NPU 组织成可训练、可推理、可调度的大规模算力系统，而不是单纯比较单卡指标。

图50：灵衢 UB 通过两级 scale up 交换组织 384 NPU + 192 CPU 协同计算域



数据来源：《超节点发展报告》，东吴证券研究所

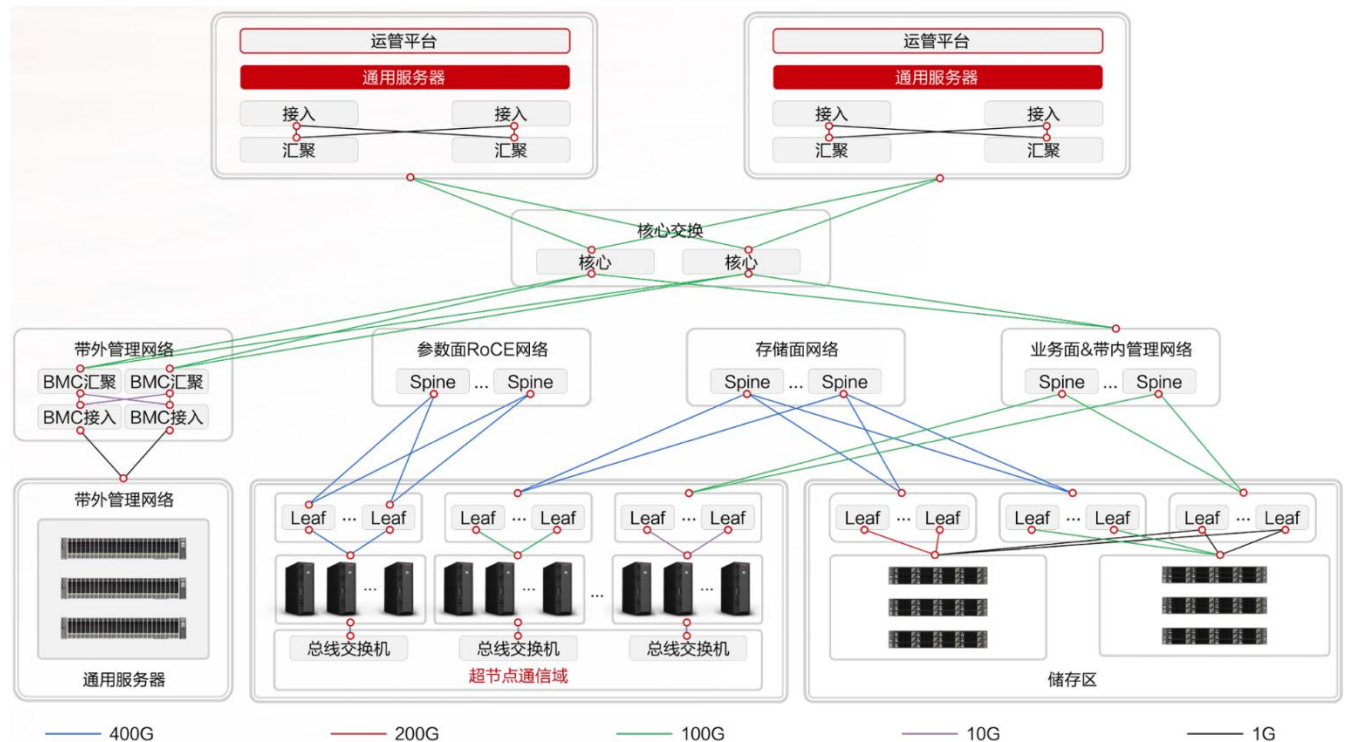
图51：灵衢协议栈



数据来源：《基于灵衢的超节点参考架构白皮书》，东吴证券研究所

华为产品演进的主线，是从 Atlas 900 A3/昇腾 384 走向 Atlas 950/960 SuperPoD，再进一步走向 SuperCluster。Atlas 900 A3 对应 384 卡级国产超节点起点，Atlas 950 SuperPoD 将单个 Scale Up 域扩大至 8192 卡，Atlas 960 SuperPoD 进一步提升至 15488 卡；SuperCluster 则把多个 SuperPoD 继续横向组织，面向 50 万卡和百万卡级集群。该路线的产业含义在于，国产算力体系正在从服务器级供给升级为超节点和超集群级供给，系统互联、资源调度和大规模稳定运行成为华为路线的核心验证点。

图52: 超节点集群组网架构 (以昇腾 384 超节点为例)



数据来源:《超节点发展报告》, 东吴证券研究所

6.3.1. Atlas 900 A3 / 昇腾 384: 国产超节点起点

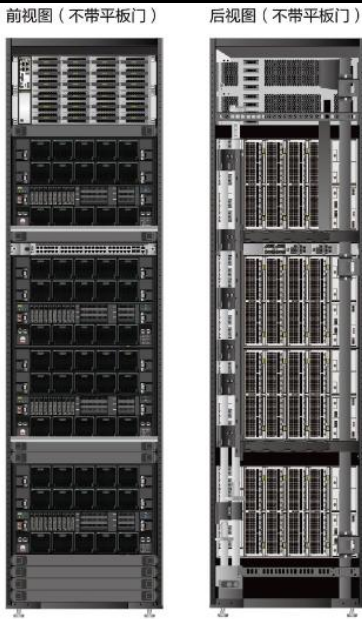
Atlas 900 A3 SuperPoD (CloudMatrix 384) 是华为国产超节点的起点产品, 其意义在于把昇腾 NPU、鲲鹏 CPU、灵衢 UB 和多柜系统首次整合为明确的 384 卡产品形态。该系统由 12 个计算柜与 4 个总线设备柜组成, 最大支持 384 张昇腾 910 和 192 颗鲲鹏 920, 通过 UB 网络把多台物理设备组织为一个高带宽协同计算域。片上内存规模最高 384×128GB, D2D 双向带宽 784GB/s, 通信单跳时延 200ns; 这些指标对应的核心变化, 是国产算力供给从服务器堆叠转向系统级协同。对华为而言, Atlas 900 A3 承担的是产品谱系中的起点锚点: 先以 384 卡验证国产超节点的互联、编址和系统交付能力, 再为后续更大规模 SuperPoD 扩展提供基础。

图53: 昇腾384将国产算力从服务器堆叠推进到多柜超节点形态



数据来源：华为计算公众号，东吴证券研究所

图55: Atlas 900 A3 机柜整机示意图



数据来源：《Atlas 900 A3 SuperPoD 超节点 技术白皮书》，东吴证券研究所

图54: Atlas 900 A3 以 384 NPU 高速互联、统一编址和百纳秒级时延构建国产 Scale Up 能力



数据来源：华为官网，东吴证券研究所

图56: Atlas 900 A3 部署落地情况（截至 2025 年 9 月 18 日）



数据来源：华为中国公众号，东吴证券研究所

图57：昇腾芯片路线持续提升算力、互联带宽与内存



数据来源：华为全联接大会 2025，邮电经济公众号，东吴证券研究所

6.3.2. Atlas 950/960 SuperPoD：国产 Scale Up 域持续扩大

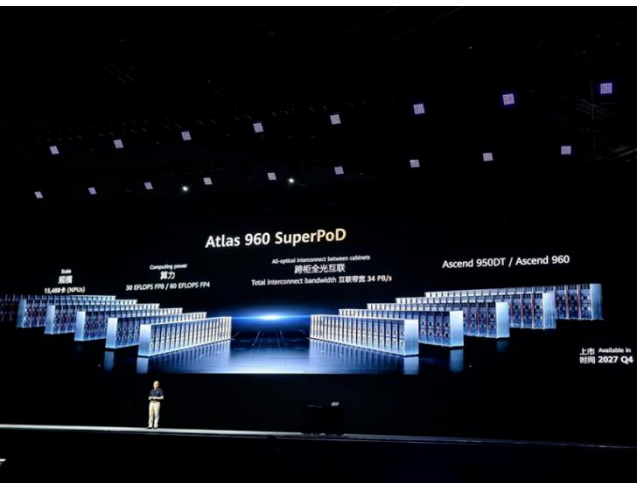
Atlas 950 SuperPoD 标志着华为超节点从 384 卡级系统进入万卡级 Scale Up 阶段。该产品以单柜 64 卡为基本扩展单元，最高支持 8192 张昇腾卡高速互联，产品形态从多柜超节点升级为 SuperPoD 级计算系统。相比 Atlas 900 A3，Atlas 950 的变化不只是规模扩大，而是将昇腾 NPU、灵衢 UB、全光互联、液冷整柜和系统软件纳入更大规模的协同计算单元，推动国产算力从单一超节点产品走向系统级供给。该产品在华为“超节点—SuperPoD—SuperCluster”路线中承担关键中间层角色，向下承接 384 卡级超节点的工程基础，向上支撑 50 万卡级 SuperCluster 的集群扩展，国产超节点由此进入更大规模系统化部署阶段。

图58：Atlas 950 SuperPoD 将采用昇腾 950DT，Scale Up 域扩展至 8,192 NPU



数据来源：华为全联接大会 2025，CSDN 技术社区公众号，东吴证券研究所

图59：Atlas 960 SuperPoD 进一步扩大至 15,488 NPU



数据来源：华为全联接大会 2025，芯果公众号，东吴证券研究所

图60: Atlas 950 SuperPoD 具备 16.3PB/s 互联带宽, 推理吞吐达 19.6mn TPS



数据来源: 华为全联接大会 2025, IT 之家, 东吴证券研究所

Atlas 960 SuperPoD 与 Atlas 950/960 SuperCluster 进一步放大了华为路线的系统属性。Atlas 960 面向 15488 张昇腾卡高速互联, 继续扩大单个 SuperPoD 的计算规模; Atlas 950/960 SuperCluster 则将多个 SuperPoD 组合为 50 万卡和百万卡级集群。华为路线由此从“做出国产超节点”推进到“组织国产超集群”, 竞争重点转向大规模部署、系统调度和生态承接能力。

6.4. 中科曙光: 开放架构下的高密度超节点系统

中科曙光的定位是国产开放超节点系统方案商, 核心能力在于以整柜系统承接多元国产算力生态。与华为更强调自有芯片和自有互联体系不同, 中科曙光的路线更偏开放适配: 通过 scaleX 高密度整柜、scaleFabric 高速网络、液冷供电和系统软件, 将不同国产加速卡、CPU、操作系统、模型框架和行业应用纳入统一基础设施。该路线的价值在于降低客户对单一算力平台的依赖, 并把公司长期积累的超算系统、服务器工程和数据中心集成能力迁移到 AI 超节点场景。

图61: scaleFabric 以软硬件开放适配承接多元国产算力生态



数据来源：中科曙光官方公众号，东吴证券研究所

6.4.1. scaleX640: 单机柜级高密度超节点

scaleX640 是中科曙光超节点产品化的核心锚点，将国产开放算力系统推进到单机柜级高密度形态。该产品定位为 640 卡 AI 超节点，双 scaleX640 可组成 1280 卡计算单元，产品重点在于将计算、互联、散热、供电和系统交付压缩进更高密度的整柜系统。相比传统多服务器集群，scaleX640 的价值在于通过高密部署和整柜协同提高单位空间内的有效算力供给，使中科曙光从服务器和超算系统供应商进一步进入高密度 AI 基础设施竞争。

scaleX640 的识别度来自“高密架构 + 正交互联 + 浸没相变液冷 + 高压直流供电”的整柜组合。单柜 640 卡要求计算、交换、冷却和供电在极高空间密度下协同，超高速正交架构压缩计算节点与交换节点之间的连接路径，浸没相变液冷承接高密度热负载，高压直流供电提升整柜电力承载能力。产品口径下，单机柜总算力超过 600PFLOPs，算力密度较同类产品最大提升 20 倍，PUE 低于 1.04；同时兼容多品牌国产加速卡和 400+ 主流大模型。上述指标体现出 scaleX640 的核心竞争力体现在在开放架构下实现高密度整柜产品化。

图62: scaleX640 以 640 卡高密整柜实现开放架构下的单柜超节点产品化



数据来源: 中科曙光官网, 中科曙光官方公众号, 东吴证券研究所

6.4.2. scaleFabric 与 scaleX 万卡超集群: 从单柜超节点走向集群级基础设施

scaleFabric 承担中科曙光超节点体系中的集群扩展角色, 使 scaleX640 从单柜高密系统继续走向多柜、万卡级部署。scaleFabric 基于 RDMA 架构, 面向 AI 训练、推理和科学计算等高性能场景, 形成从网卡、交换机到管理软件的高速互连方案。中科曙光通过自研高速网络把系统能力从整柜内部延伸到集群层面, 把高密度整柜进一步组织为可扩展的国产智算基础设施。

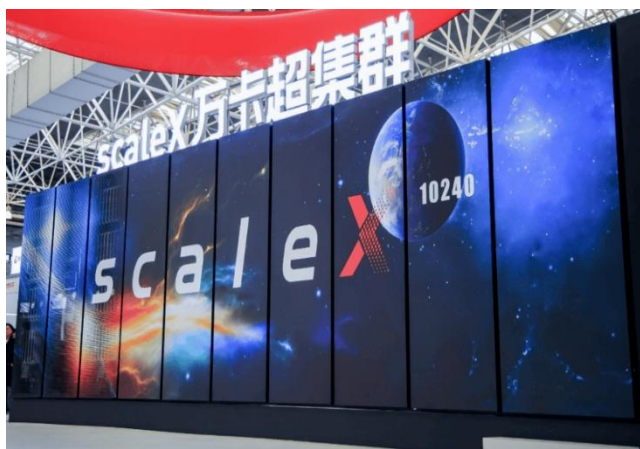
图63: scaleFabric 构建从网卡、交换机到管理软件的高速互连底座



数据来源: 中科曙光官方公众号, 东吴证券研究所

scaleX 万卡超集群体现了中科曙光从单柜超节点向集群级系统扩展的产品路径。该系统由 16 个 scaleX640 超节点通过 scaleFabric 高速网络互联，可实现 10240 块 AI 加速卡部署，总算力规模超过 5EFlops。scaleX640 提供单柜高密度算力，scaleFabric 负责超节点间连接，数字孪生与智能调度提升故障定位和系统管理效率。中科曙光路线由此形成“单柜高密度超节点—高速互连网络—万卡超集群”的产品层级，开放高密度系统开始从单点产品走向集群级基础设施。

图64: scaleX 万卡超集群标志中科曙光从单柜超节点走向集群级基础设施



数据来源：中科曙光官方公众号，东吴证券研究所

图65: scaleX 万卡超集群由 16 个 scaleX640 互联组成，形成 10240 卡、超 5EFLOPS 系统能力

scaleX万卡超集群

面向万亿参数大模型&科学智能复杂任务场景

1套万卡规模超集群系统
由多个scaleX640超节点
通过scaleFabric高速网络互连而成

10240卡

总卡数

>5EFlops

总算力

>650TB

HBM总容量

>18PB/s

HBM总带宽

>4.5PB/s

片间互连总带宽

>500TB/s

柜间互连总带宽

数据来源：中科曙光官方公众号，东吴证券研究所

6.5. 云厂商与开放互联：真实负载牵引架构创新

云厂商与开放互联路线的核心，是从真实负载和生态开放出发探索更灵活的 **Scale Up 方案**。阿里更侧重云侧资源池化和推理效率优化，字节、腾讯更侧重以太网 Scale Up 协议探索，ETH+ 则进一步把国内开放互联推进到芯片、样机和生态导入阶段。该路线拓展了超节点架构的演进方向，使系统设计从封闭高性能互联延伸到云厂商负载牵引、以太网产业链复用和跨厂商协同。

6.5.1. 阿里磐久 AL128：云侧资源池化牵引超节点重构

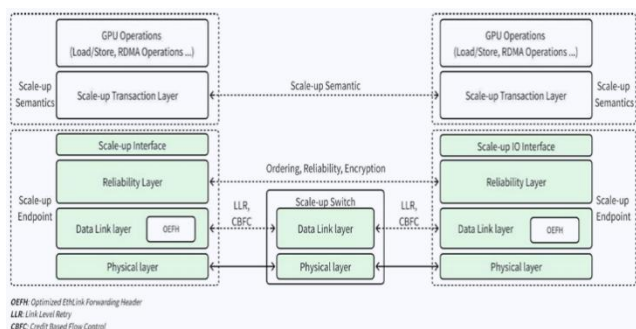
阿里磐久 AL128 体现了云厂商从真实训练和推理负载出发重构超节点架构的路径。该系统以 128 卡为核心规模，通过 ALink、单级交换、GPU 节点与交换节点正交互联，以及 CPU/GPU/交换节点解耦，缩短超节点内部通信路径，并提升资源组织弹性。Prefill、Decode、长上下文、多模态和高并发推理对算力、缓存、网络和调度提出不同要求，AL128 通过资源池化和 KV Cache 优化，使计算、主控、交换和内存资源能够围绕

业务负载重新组合。其产品价值集中在云侧推理效率、资源利用率和单位算力产出提升。

6.5.2. 字节 EthLink 与腾讯 ETH-X: 以太网 Scale Up 的协议探索

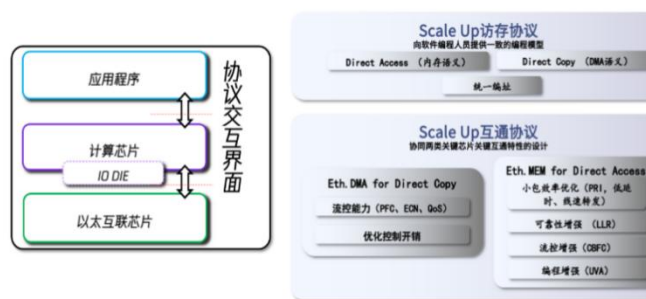
字节 EthLink 与腾讯 ETH-X 代表国内云厂商推动以太网进入 Scale Up 场景的两类探索。EthLink 面向大规模 AI 集群负载,对原生以太网进行 GPU Scale Up 优化,使其同时支持低时延小块同步数据和高带宽大块异步传输,推动以太网从传统节点间网络进入节点内高带宽互联。ETH-X 更强调开放协议和原型验证,围绕 GPU-GPU 访存、GPU-Switch 互通和链路可靠性建立接口,连接芯片、整机、连接器/线缆和云侧需求。两条路线的共同价值,在于用以太网产业基础降低超节点互联的生态门槛。

图66: 字节 EthLink 协议栈



数据来源: 字节跳动《GPU Scale up 互联技术白皮书》, 东吴证券研究所

图67: ETH-X Scale U 互联协议示意图

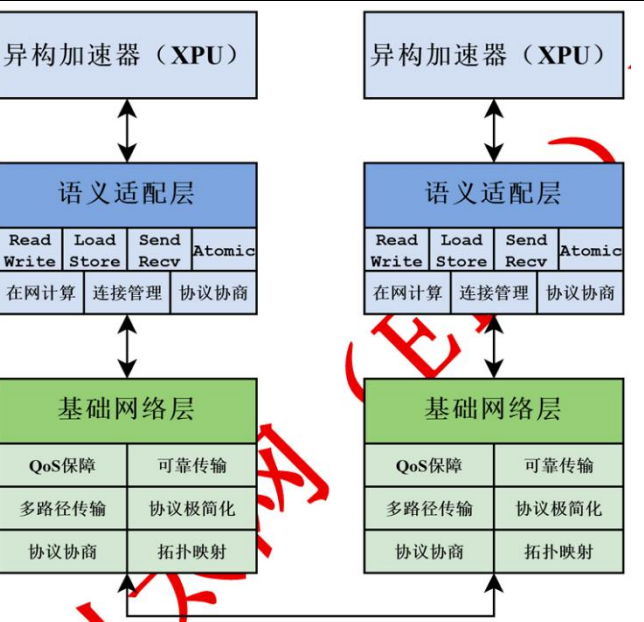


数据来源: ODCC, 东吴证券研究所

6.5.3. ETH+, ERack+ 与 UPN 512: 国内开放超节点从协议走向样机

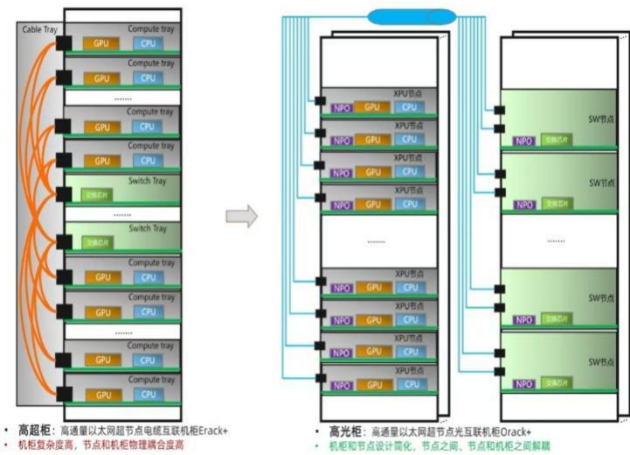
ETH+ 体系代表国内开放互联路线从协议讨论进入芯片、样机和生态导入阶段。该体系围绕 Scale Up 和 Scale Out 场景推进,覆盖国产 400G 智能网卡芯片、25.6T 交换芯片、硅光芯片、ERack+ 64 超节点和 UPN 512 光互联方案,产业意义在于通过联盟方式聚合云厂商、芯片厂商、系统厂商和科研机构,形成国产开放互联的共同接口。与单一厂商闭环路线相比,ETH+ 更强调开放生态和部件协同,后续若在芯片量产、整柜验证和客户场景中持续推进,有望成为国产超节点的重要补充路径。

图68: ETH+ Scale up 协议架构



数据来源: 高通量以太网联盟《Scale-Up 场景互连协议白皮书》, 东吴证券研究所

图69: 基于 ETH+光互联的超节点解决方案



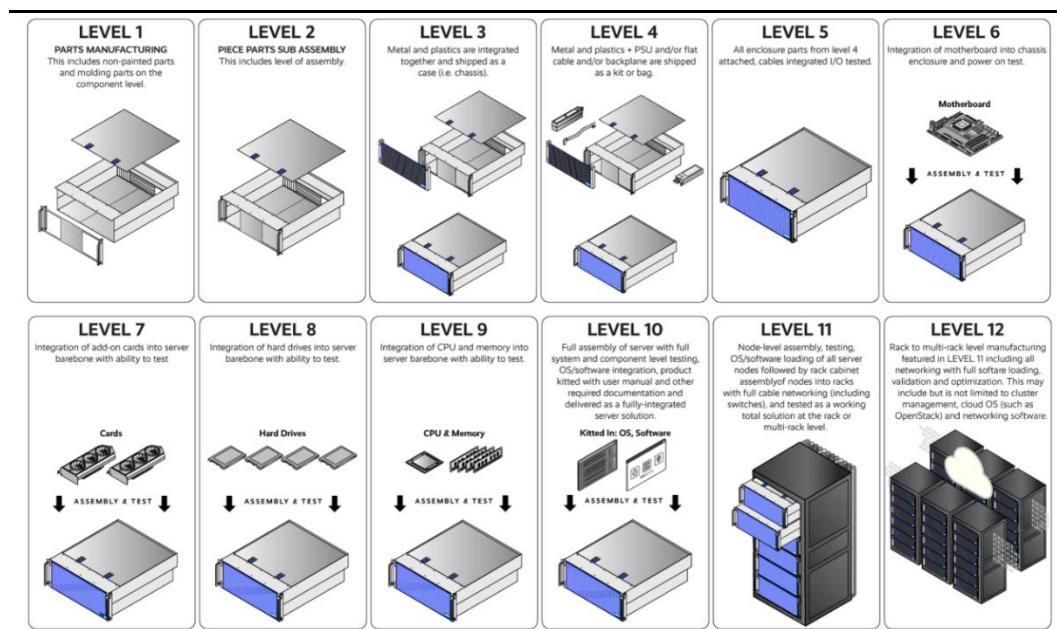
数据来源: CCF HPC China 2025 大会, 云深知网络公众号, 东吴证券研究所

7. 超节点重塑服务器厂商价值链：整柜交付与客户生态成为分化核心

7.1. 超节点推动服务器厂商角色升级：从整机制造走向整柜系统交付

AI 服务器需求正在从通用计算向加速计算快速迁移。2025 年全球服务器市场收入同比增长 80.4%；2025Q4 x86 服务器收入 698 亿美元，同比增长 16.9%，非 x86 服务器收入 555 亿美元，同比增长 146.4%，增速显著分化；同期搭载 GPU 的服务器收入同比增长 59.1%，占全球服务器市场收入超过 50%，AI 算力投入推动服务器市场加速向异构计算和高密度算力系统迁移。超节点在这一背景下推动服务器厂商交付单位由单机服务器上移至整柜系统，客户采购对象从单台设备转向经过整柜级集成、验证和部署的高密度算力单元，服务器厂商能力边界也由单机设计和装配延伸至机柜级系统工程。

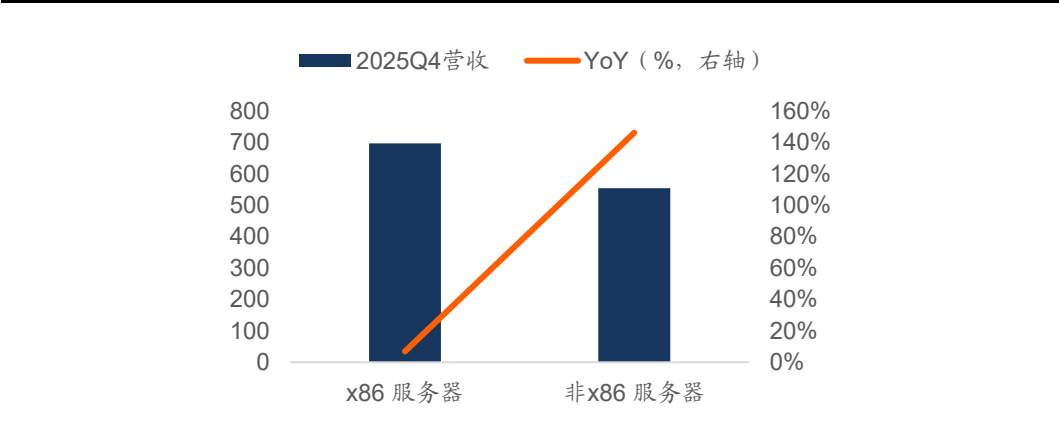
图70：从 L10 到 L11 整机柜、L12 多机柜交付



数据来源：AMAX，东吴证券研究所

交付层级上移是服务器厂商价值重估的核心。传统 AI 服务器更多停留在 L10 层级的单机装配与测试，超节点则推动厂商向 L11 整机柜交付、甚至 L12 多机柜交付延伸，参与环节从主板、机箱、电源和散热扩展到 Compute Tray、Switch Tray、液冷、集中供电、整柜烧机、链路验证和现场部署。价值留存不再只来自整机制造，而是来自托盘级部件、整柜集成、系统测试和项目交付能力，服务器厂商的竞争门槛也随之从制造效率提升到系统工程能力。

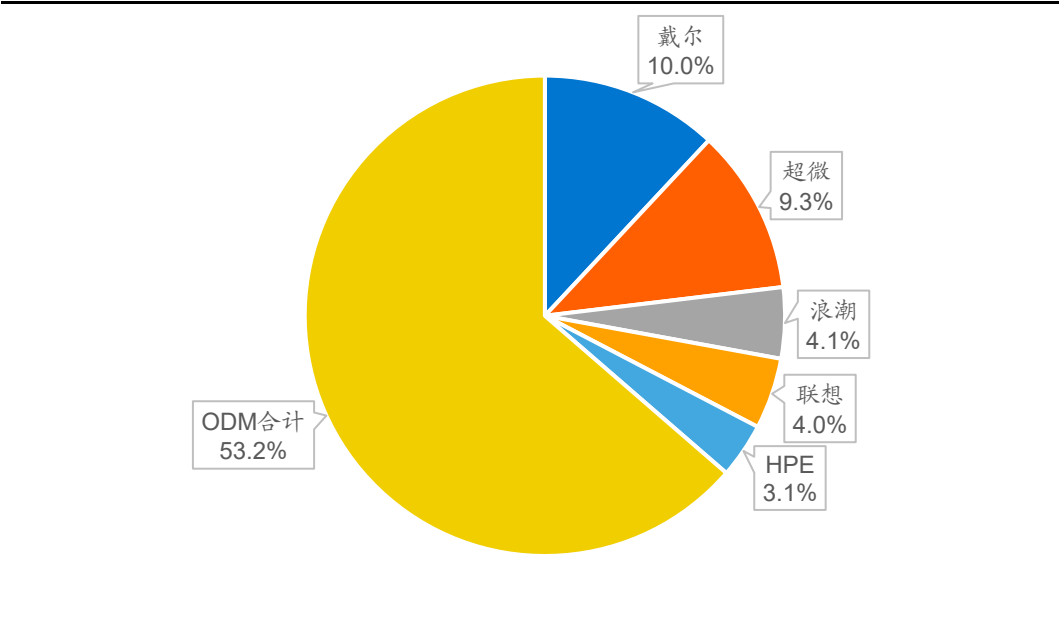
图71：2025Q4 分架构营收（亿美元），非 x86 狂飙



数据来源：IDC，东吴证券研究所

客户生态和项目交付能力正在成为服务器厂商分化的核心变量。根据 IDC，2025Q4 ODM Direct 占全球服务器收入 53.2%，同比增长 60.5%，已显著高于任一单一 OEM 厂商，反映头部云厂商和 AI 基础设施客户对定制化整柜交付、快速迭代和供应链协同的需求提升。超节点高度依赖芯片平台、互联协议、系统软件和客户数据中心条件，服务器厂商能否进入主流云厂商、运营商和国产算力生态，决定其能否参与早期适配、样机验证和后续批量交付。随着超节点从样机验证走向规模部署，客户绑定深度、整柜工程能力和项目兑现节奏将成为服务器厂商份额分化的关键。

图72：2025Q4 全球厂商份额，ODM 首次超过 OEM 总和



数据来源：IDC，东吴证券研究所

7.2. 工业富联：全球 AI 服务器与整柜制造链核心平台

依托鸿海体系的全球制造网络,工业富联是全球 AI 算力基础设施从服务器整机走向整柜系统时的核心制造平台。公司业务覆盖云计算、通信网络设备、精密机构件和工业互联网,客户认证、供应链组织和规模化制造能力构成其参与全球 AI 基础设施建设的基础。超节点时代,制造平台的能力不再停留在服务器整机装配,而是延伸到高密度装配、液冷导入、供电集成、整柜测试和跨区域交付。工业富联由此从传统电子制造平台,进一步转向全球 AI 算力基础设施制造链的核心参与者。

英伟达平台从 GB200/GB300 向下一代系统演进,正在放大工业富联的整柜制造价值。GB200/GB300 NVL72 将 AI 服务器形态推向机柜级系统,计算、交换、散热、供电等模块需在整柜内完成集成、装配与系统级验证。工业富联作为英伟达 GB200/GB300 机柜级系统的重要 ODM 厂商,增长驱动正由传统 AI 服务器出货扩张进一步转向高价值量整柜系统放量。随着英伟达平台继续向 Vera Rubin 等下一代架构演进,整柜系统的功耗密度、互联复杂度和系统集成要求进一步提升,公司能否持续承接下一代高价值量机柜产品,将影响其在英伟达 AI 基础设施供应链中的份额和盈利弹性。

高速网络设备强化了工业富联在 AI 集群制造链中的参与深度。随着 AI 集群由单柜向 POD 和 AI Factory 演进,客户投入由计算机柜进一步延伸至高速网络设备,交换机、CPO 和高带宽网络成为整柜系统之外的重要增量。公司交换机业务受集群规模提升(Scale-up)与节点扩展(Scale-out)双重驱动,800G 产品进入快速放量阶段,根据公司披露,仍将是 2026 年出货主力;同时,公司已与多家客户联合开发 CPO 与 1.6T 交换机并取得在手订单,CPO 全光交换机样机已开始出货,下一代 Vera Rubin NVL144 平台交换机亦已完成开发。公司同时覆盖计算侧整柜制造和网络侧设备制造,有助于提升其在客户数据中心建设中的参与层级和单客户价值量。

工业富联的壁垒来自全球制造、客户认证和大规模交付能力,约束则来自制造属性和客户集中度。AI 整柜系统对制造平台的要求进一步提升,产线自动化、良率爬坡、供应链组织和跨区域交付共同影响订单兑现能力。2026Q1 公司营业收入同比增长 56.5%,归母净利润同比增长 102.6%,公司表示主要受益于 AI 算力需求增长、主要客户份额提升以及产品结构持续优化;经营活动现金流净额同比大幅改善,主要受销售回款增加驱动。2026H1 GPU AI 机柜、ASIC AI 机柜出货量分别同比增长 3.2 倍和 3 倍,显示公司在 AI 整柜业务中的规模交付能力持续兑现。与此同时,2025 年前五大客户销售额占比达 62.0%,叠加制造环节较高的材料成本占比,客户集中度、产品毛利率及平台迭代周期仍是影响盈利弹性的主要约束。

图73：2026 一季度 AI GPU 机柜、AI ASIC 服务器与 800G 交换机高增



数据来源：工业富联公众号，东吴证券研究所

7.3. 华勤技术：平台型 ODM 向 AI 服务器和数据中心产品延伸

从多品类智能硬件 ODM 走向数据中心，是华勤技术在 AI 基础设施周期中的主要扩张方向。公司长期深耕智能手机、笔电、平板、AIoT 等产品，形成了研发设计、供应链管理、规模制造和全球交付的综合能力。进入 AI 基础设施周期后，公司围绕 AI 服务器、通用服务器、存储服务器和交换机拓展数据中心产品，并与国内头部云厂商保持较深合作。相较于深度绑定全球闭环整柜生态的制造平台，华勤更偏向平台型 ODM 向算力基础设施外延扩张，弹性来自数据中心业务放量和国产算力生态扩容。

图74：“3+N+3”产品平台覆盖多元终端与服务器，平台型 ODM 能力向数据中心延伸



数据来源：华勤技术官网，东吴证券研究所

AI 服务器业务是华勤进入数据中心平台的核心抓手。公司数据中心业务自 2017 年

起布局 AI 服务器、通用服务器、存储服务器及交换机等产品，并持续提升在头部 CSP 客户中的采购份额。2025 年下半年以来，公司数据中心业务加速向国产 GPU 平台切换，国产 AI 服务器业务快速起量；2026Q1 数据中心业务虽受去年 H20 高基数影响同比略有下降，但环比延续增长，超节点产品已于 Q2 开始小批量发货，公司预计下半年进入批量交付阶段、全年收入超过 100 亿元。AI 服务器业务由单机产品逐步向国产 GPU 适配、超节点及全栈数据中心方案延伸，正成为公司向高价值算力基础设施升级的重要通道。

图75：数据中心业务持续高增，AI 服务器与超节点加速产品升级

数据中心

收入近翻倍增长，市场份额持续提升，**超节点**技术行业领先



数据来源：华勤技术公众号，东吴证券研究所

超节点业务提升了华勤数据中心产品的工程复杂度和价值弹性。公司已推出面向下一代智算中心的 AI 超节点方案，并在计算节点、网络节点、整机架构、高速互联、液冷和供电等环节形成布局，数据中心业务正在从单机服务器向整柜级系统能力推进。ETH-X 开放超节点原型机在华勤东莞智能制造基地完成下线点亮，也体现其在开放互联生态和整柜制造验证中的参与位置。公司预计超节点项目将从小批量发货走向规模量产，若交付节奏兑现，AI 业务有望从服务器出货扩张进一步走向整柜系统价值提升。

图76: ETH-X 原型机在华勤技术东莞智能制造基地点亮



数据来源: ODCC, 东吴证券研究所

华勤的竞争优势在于平台化 ODM 能力、国内 CSP 客户基础和产品线延展, 高端整柜交付仍处于爬坡验证阶段。公司具备客户协同、快速研发、供应链组织和多基地制造能力, 有助于承接云厂商定制化需求, 并在国产算力多平台并行环境下提升适配效率。数据中心业务仍处于快速扩张期, AI 服务器、交换机和超节点业务均具备收入弹性; 同时, 液冷整柜、超节点规模交付、长期运行验证和客户份额提升仍需要通过后续订单和项目交付持续验证。

7.4. 浪潮信息: 国内 AI 服务器龙头, 液冷整柜与智算中心交付能力突出

作为国内服务器龙头, 浪潮信息的价值重心正在从 AI 服务器规模供给延伸到智算中心系统方案。公司长期深耕服务器、存储、网络 and 算力基础设施, 产品覆盖通用计算、AI 计算、边缘计算、关键计算和开放网络等场景, 客户基础横跨互联网、运营商、政企、金融和行业市场。相较于以制造交付见长的平台型 ODM 厂商, 浪潮信息的优势更集中在服务器产品定义、软硬件适配、行业客户覆盖和国产算力生态承接。超节点时代, 公司价值不只来自 AI 服务器出货规模, 也来自其能否把计算、存储、网络、液冷和系统软件整合为更高层级的智算基础设施方案。

AI 服务器业务构成浪潮信息向超节点升级的能力基础。公司在 AI 服务器领域长期投入, 已形成训练服务器、推理服务器、AI 存储、网络互联和管理软件的产品组合,

并通过联合设计开发模式与下游客户保持较深协同。生成式 AI 训练和推理对多机多卡系统的通信效率、资源调度和长期稳定运行提出更高要求，单台服务器性能已经难以单独决定客户体验。浪潮信息在国产及海外芯片平台适配、液冷服务器和 AI 管理软件上的积累，使其具备从单机产品供应向系统级方案交付延伸的基础。同时，公司持续布局融合存储、超级 AI 以太网及 AIStation 等管理平台，产品能力已覆盖计算、存储、互连和管理等智算基础设施核心环节。

元脑 SD200 是浪潮信息从 AI 服务器向国产超节点系统延伸的关键产品。该产品面向国产 AI 芯片高速互连和大模型训推场景，在单机内实现 64 路国产 AI 芯片高速统一互连，并通过低延迟内存语义通信提升多芯片协作效率，支持单机承载 4 万亿参数单体模型，或部署多个万亿参数模型组成智能体应用。同期发布的元脑 HC1000 则更偏向超扩展推理系统，基于 DirectCom 架构聚合海量国产 AI 芯片，强化大规模推理吞吐和 token 成本优化。SD200 与 HC1000 分别对应国产超节点互联和推理侧规模扩展，浪潮信息的产品边界由此从传统 AI 服务器扩展到更高密度的国产计算单元，服务器、互联、系统管理和可靠性验证被纳入同一产品能力框架。

图 77：元脑 SD200 与 HC1000 分别强化国产芯片互联和大规模推理能力



数据来源：2025 人工智能计算大会，东吴证券研究所

浪潮信息的后续分化取决于国产超节点放量、客户结构改善和系统级交付兑现。公司具备国内服务器龙头的产品覆盖、客户基础和项目经验，AI 服务器需求高景气为其提供规模支撑，超节点则为其提升系统价值和盈利弹性提供新抓手。约束在于，互联网大客户采购占比提升可能压缩议价空间，高端芯片供应和客户资本开支节奏仍会影响订单兑现。若 SD200 等产品形成批量交付，公司有望从国内 AI 服务器龙头进一步升级为国产超节点系统的重要供给方。

7.5. 紫光股份（新华三）：网络底座与政企智算优势强化系统交付能力

以新华三为核心业务平台，紫光股份在 AI 算力基础设施中的特色来自网络底座与政企 ICT 解决方案能力。公司业务覆盖计算、网络、存储、云、安全和终端，客户基础横跨政企、运营商、金融、互联网和海外市场，具备较完整的 ICT 基础设施产品与解决方案能力。超节点从单柜计算单元继续走向多柜集群和智算中心后，网络连接、资源调度、系统管理和行业场景适配的重要性提升。新华三的价值不只在 AI 服务器出货上，更应放在计算、网络、云平台 and 政企方案协同形成的系统交付能力上。

网络能力是新华三区别于一般服务器厂商的核心抓手。AI 集群规模扩大后，训练和推理效率越来越依赖 Scale Out 网络质量，数据中心交换、无损以太、低时延转发、拥塞控制和统一运维会直接影响算力释放效率。新华三在以太网交换机、数据中心交换机、路由器和 WLAN 等领域长期保持较强市场位置，并围绕 800G、国芯智算交换机、CPO 硅光交换机和智算网络持续布局。相比只具备服务器整机能力的厂商，公司能够同时参与计算节点和网络底座建设，在智算中心项目中更容易形成组合交付优势。

图78：新华三 102.4T 智算交换机通过光电并行提升密度、吞吐与低时延能力



数据来源：NAVIGATE 2026，东吴证券研究所

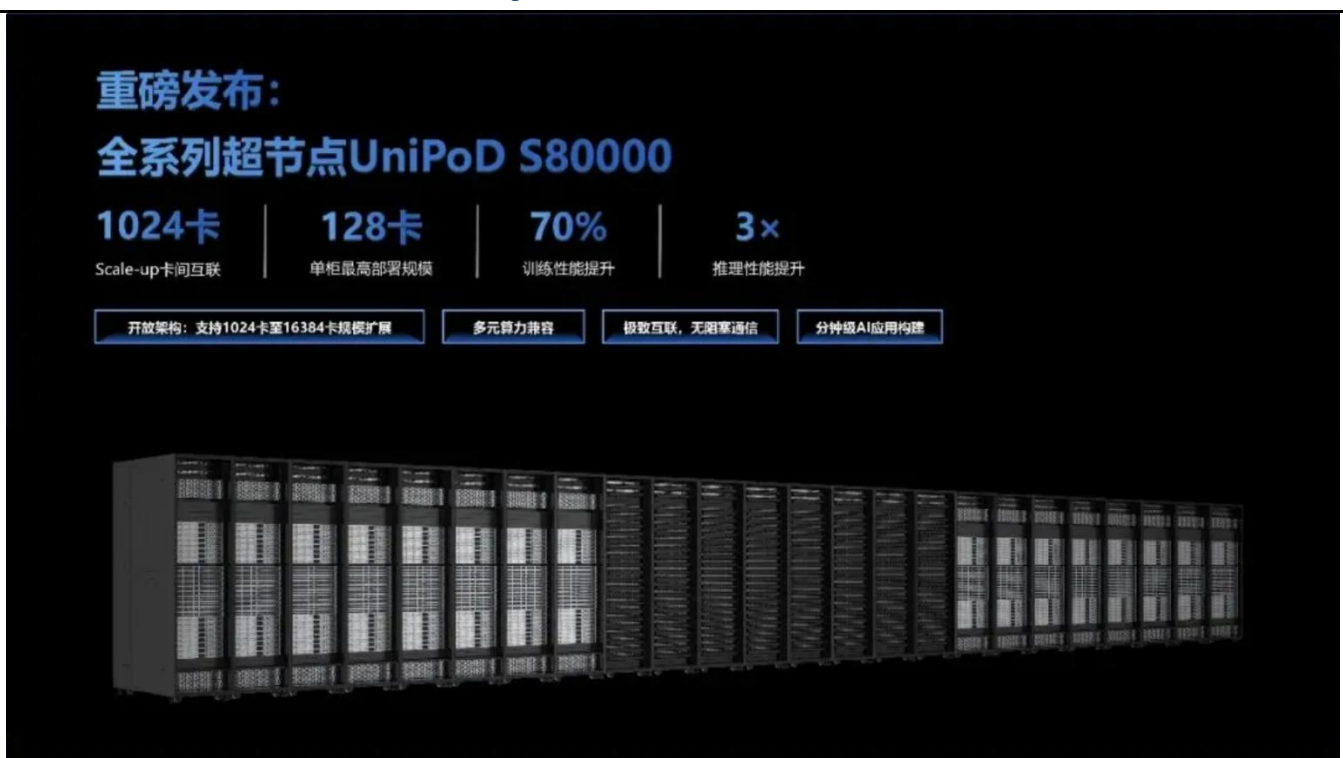
图79：新华三智算广域网融合核心路由器与 800G 交换机，支撑千公里级无损传输



数据来源：NAVIGATE 2026，东吴证券研究所

服务器与超节点产品补齐了新华三在算力侧的系统入口。新华三在 X86 服务器、GPU 服务器、存储和计算平台上具备较完整产品基础，能够覆盖通用计算、AI 训练推理、存储和边缘计算等场景。UniPoD S80000 超节点进一步将公司能力从服务器和网络设备延伸到高密度 AI 计算系统，单柜可部署 32、64、最高 128 张 AI 加速卡，柜内可实现卡间全互联通信，互联带宽较传统 8 卡服务器组网提升 8 倍；针对 MoE 模型，单柜训练效率最大提升 10 倍，单卡推理效率最大提升 13 倍。其产业意义在于，新华三可以把传统强项网络能力嵌入超节点和智算中心架构中，形成区别于单纯服务器厂商的系统方案能力。

图80: UniPoD S80000 支持 1024 卡 Scale Up 互联与 128 卡单柜部署



数据来源: NAVIGATE 2026, 东吴证券研究所

政企与运营商客户基础提升了新华三在国产智算项目中的落地能力。国内超节点和智算中心建设并不只来自互联网大模型客户，政企、运营商、金融、能源、教育和医疗等行业同样需要算力基础设施、数据中心网络、云平台和安全能力的组合交付。新华三在这些客户中的长期积累，使其更适合承接行业智算中心、国产化替代和混合云场景下的系统建设需求。若超节点产品规模部署和国产智算项目持续兑现，新华三有望在国内服务器厂商分化中占据“网络底座 + 服务器 + 政企解决方案”的特色位置。

图81：全栈协同定义新一代 AI 基础设施



数据来源：NAVIGATE 2026，东吴证券研究所

7.6. 超聚变：运营商与政企场景中的整机柜液冷系统交付平台

承接原华为 X86 服务器业务基础，超聚变更适合在运营商、政企和行业智算中心的国产服务器体系中观察。公司产品覆盖通用服务器、AI 服务器、操作系统、管理软件和算力基础设施解决方案，客户主要分布在运营商、政企、金融、互联网和行业数据中心等场景。相比更偏 ODM 制造的平台型厂商，超聚变的特点在于国产服务器体系、整机柜液冷产品和行业客户交付经验。公司上市辅导推进后，业务规模、客户结构和治理能力的市场可见度有望提升。

图82：超聚变产品体系覆盖服务器、操作系统与 AI 解决方案，强化国产智算基础设施交付能力

 数据中心	 边缘计算	 操作系统
FusionServer通用服务器 >	FusionXpark™ 随身智能体开发平台 >	FusionOS服务器操作系统 >
FusionServer AI服务器 >	工作站X3 8000 >	
FusionPoD整机柜液冷服务器 >	工控机 8588 >	
FusionDirector服务器管理软件 >		
FusionWatt服务器电源 >		
 AI解决方案	 超融合解决方案	 高性能解决方案
Fusion AI Space大模型加速引擎 >	FusionOne HCI解决方案 >	FusionOne HPC解决方案 >
FusionOne AI场景化解决方案 >	FusionOne for VMware解决方案 >	
	FusionOne for Microsoft解决方案 >	

数据来源：超聚变官网，东吴证券研究所

FusionPoD for AI 是超聚变面向高密度 AI 算力场景推出的整机柜液冷服务器，集中体现其整柜工程能力。该产品面向大规模、超大规模模型训练和推理场景，将服务器节点、液冷散热、集中供电、机柜结构和运维管理纳入整柜级设计，支持 105—240kW 集中供电、单柜 64—128 加速卡和多样性算力兼容。其产品价值集中在高功率密度下的整柜液冷、供电集成、快速部署和无人智检运维，能够对应运营商和政企客户对智算中心建设效率、长期运行稳定性和规模化运维的要求。

图83：FusionPoD for AI 集成液冷、供电与运维管理，支撑 64—128 卡整机柜交付



数据来源：超聚变官网，东吴证券研究所

运营商智算项目验证了超聚变整机柜液冷系统的规模部署和场景适配能力。杭州电信智算中心规划建设万卡规模国产算力集群，FusionPoD for AI 通过架构开放和多算力兼容支撑双生态液冷智算集群，并与天翼云“云骁”“息壤”等平台协同，实现多卡多机调度、算力加速和跨网算力分发。中国联通互联网应用创新基地项目则突出其全液冷一体化和无人运维能力，系统实现计算、存储、网络、电源全节点液冷覆盖，制冷 PUE < 1.05，单机柜供电突破 120kW，部署效率提升 10 倍以上，运维成本降低 50% 以上。两类项目共同说明，超聚变的竞争力已经从单一服务器产品延伸至运营商级智算中心的液冷整柜交付、国产算力兼容和长期运维场景。

图84：静音全液冷一体化系统联合创新样板点



数据来源：超聚变官网，东吴证券研究所

7.7. 宁畅：以液冷整柜与高密 AI 服务器切入智算场景，超节点卡位仍待规模验证

宁畅的差异化在于以高密服务器、液冷产品和 JDM 定制能力切入中大型智算场景。公司长期聚焦服务器研发、生产、部署和运维，产品覆盖通用机架服务器、AI 服务器、多节点服务器、边缘计算服务器和 JDM 定制化方案，客户场景面向互联网、电信、金融、科研、教育和数据中心等行业。相较于综合型 ICT 厂商，宁畅更强调高密计算、液冷服务器、快速交付和场景化定制，适合在中大型智算项目中观察其整柜工程和灵活部署能力。

B8000 整机柜系统是宁畅从服务器单机供给走向整柜系统交付的核心产品。该产品采用 42U 整机柜形态，支持不低于 32U 节点部署，2U 节点可搭载双路第五代英特尔至强可扩展处理器，面向云计算、虚拟化、大数据和高性能计算等高密度场景。B8000 将计算节点、交换、电源箱、冷却单元和机柜平台统一交付，通过电、网、液三总线盲插降低现场部署和运维复杂度。液冷侧，全冷板覆盖率可达 90% 以上，核心部件降温超过 15℃，PUE 可低至 1.09；供电侧，集中供电与冗余设计提升整柜供电密度和可维护性。该产品体现宁畅在整柜液冷、集中供电和快速部署环节的能力延伸，而非单台服务器参数竞争。

图85：宁畅 B8000 以 42U 整机柜形态整合计算、交换、供电与液冷单元



数据来源：宁畅官网，东吴证券研究所

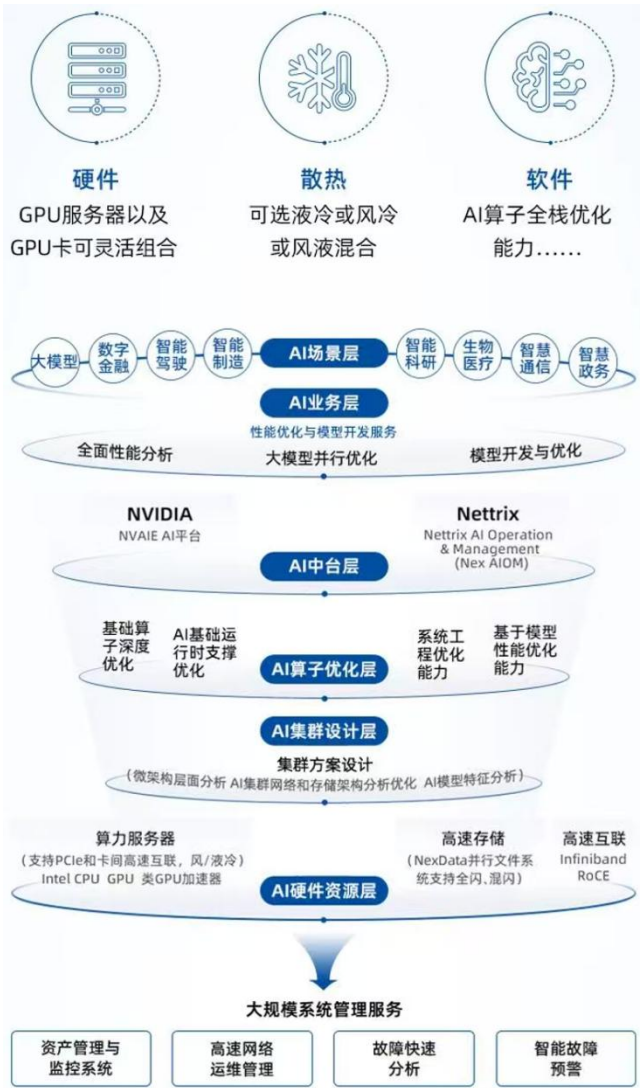
AI 服务器和多节点产品矩阵强化了宁畅向高密智算场景延伸的基础。X680 G55 是公司当前高端 AI 服务器的代表型号，采用 8U 形态，支持第四代/第五代英特尔至强处理器和 8 颗 NVLink GPU，面向 AI 大模型集群训练和大规模部署；X640 G50 则以 4U 形态支持最多 10 张双宽 GPU，补充高密推理和通用 GPU 计算场景。结合 B8000 整机柜系统，宁畅的产品能力正在从单机 AI 服务器延伸到高密整柜液冷和大模型一体化交付。其超节点卡位仍处于规模验证阶段，后续弹性取决于 B8000 等整柜产品在数据中心项目中的复制节奏，以及 AI 服务器产品与客户训推负载需求的绑定深度。

图86: DeepSeek 大模型一体机解决方案产品矩阵



数据来源: 宁畅公众号, 东吴证券研究所

图87: 宁畅“全栈全液”AI基础设施方案架构



数据来源: 宁畅公众号, 东吴证券研究所

8. 系统仿真与设计验证: 从工程复杂度走向可复制交付

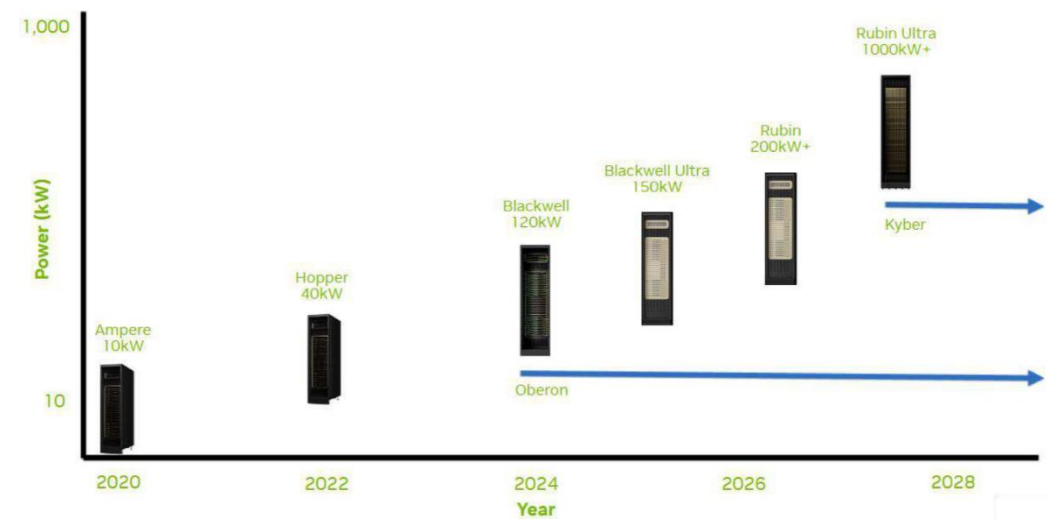
8.1. 单柜功耗跃迁推动设计验证前置

超节点的工程复杂度首先来自单柜功耗快速上行。AI 服务器从 Ampere、Hopper 时代进入 Blackwell、Rubin 时代后, 整柜功率从十千瓦级、数十千瓦级跃升至百千瓦级, 后续多机架高带宽域还会进一步推高单元功耗。功率密度变化会同步放大散热、供电、结构承载、线缆布局和现场部署约束, 整柜产品的验证难度已经明显高于传统服务器。

验证前置成为高密整柜产品化的必要环节。样机跑通只能证明架构可行, 规模交付还需要在设计阶段提前锁定热负荷、供电冗余、流量分配、管路压降、连接可靠性和机

房侧条件。后期返工不再只是单个部件修正，而会牵动冷板、Manifold、快接头、Power Shelf、Busbar、线缆路径和客户现场部署方案。具备前置仿真和场景验证能力的厂商，更容易缩短客户认证周期，并降低整柜系统批量交付风险。

图88: AI 整柜功耗从十千瓦级跃升至百千瓦级，并向兆瓦级系统演进



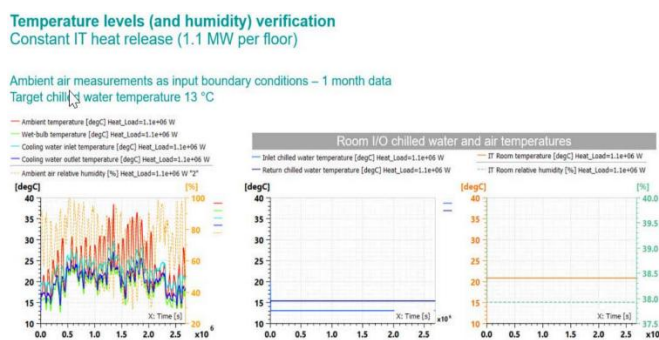
数据来源：英伟达官网，东吴证券研究所

8.2. 热管理仿真与虚拟调试决定液冷方案收敛效率

液冷系统的设计难点，在于整柜热负载、设施侧冷源和控制策略需要同步收敛。高密计算柜内部热源集中，负载波动持续，冷却塔、冷水机组、水侧节能器、泵、CRAH/CDU与机房环境之间存在联动关系。系统级仿真和虚拟调试可以把 IT 负载、冷冻水温度、流量分配、环境边界和控制逻辑纳入同一模型，在物理部署前验证冷却系统的运行边界。

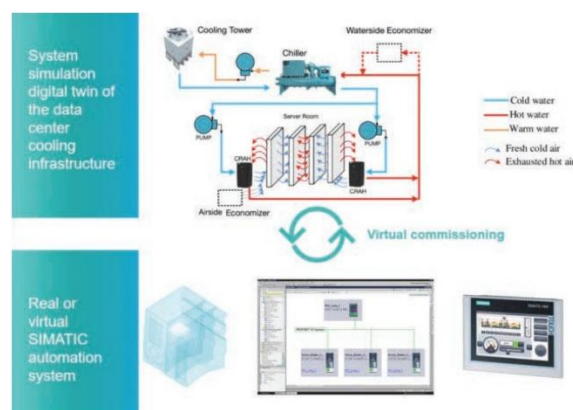
温湿度、冷却水进出口温度和机房环境稳定性，是液冷方案能否进入工程交付的重要验证对象。高功率整柜的风险不只来自单点芯片温度，也来自负载持续释放后的冷源响应、管路压降、流量分配和控制策略波动。通过仿真结果与目标工况对照，厂商可以在样机前识别局部热点、冷量不足和控制逻辑异常，减少后续改版和现场调试成本。热管理能力由此从散热部件设计，升级为冷却系统建模、虚拟调试和客户现场适配能力。

图89: 数据中心温湿度与冷冻水工况仿真



数据来源: Siemens 官网, 东吴证券研究所

图90: 冷却系统虚拟调试与控制策略验证



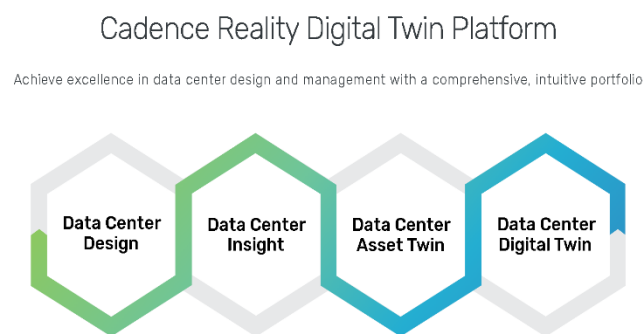
数据来源: Siemens 官网, 东吴证券研究所

8.3. 数字孪生与多物理场仿真推动整柜设计从经验走向模型驱动

超节点整柜不是单一硬件产品,而是计算、网络、供电、冷却、结构和运维系统的组合体。数字孪生平台的价值,在于把数据中心设计、运行洞察、资产管理和系统仿真纳入统一模型,使设计团队能够在物理交付前评估布局变化、功率分配、冷却效率和运行风险。随着高密整柜从单柜走向多柜和集群,验证对象也从服务器内部扩展到机柜、机房和设施侧基础条件。

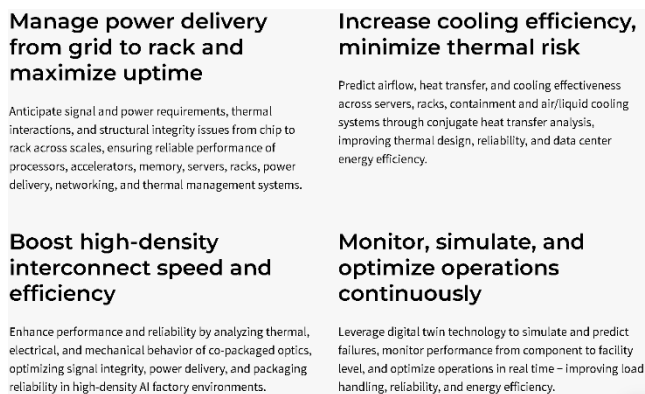
多物理场仿真进一步扩大了设计验证的边界。供电路径、热传导、气流/液流、结构可靠性、高速互联和运维状态之间高度耦合,单一维度优化容易把风险转移到其他系统。面向 AI Factory 的仿真能力,需要同时覆盖从芯片、加速器、内存、服务器到机柜和机房的多层级约束,并在运行阶段持续监测、预测和优化。对服务器厂商而言,数字孪生和多物理场仿真不是独立工具,而是把高密整柜设计、测试、部署和运维经验沉淀为可复制工程能力的基础。

图91: 数字孪生平台贯通设计、运行、资产与仿真



数据来源: Cadence 官网, 东吴证券研究所

图92: AI Factory 多物理场仿真体系



数据来源: Ansys 官网, 东吴证券研究所

8.4. 整柜测试与交付流程决定规模复制能力

整柜交付的核心，是把客户最终接收的系统状态提前纳入验证范围。超节点出厂前需要完成节点装配、整柜集成、网络与供电连接、线缆标识、满载烧机、互操作测试和性能报告；到场后还要完成软件部署、驱动配置、液冷接入、部署测试和运维支持。测试边界越接近真实训练、推理和长期高负载场景，客户现场的不确定性越低。

规模交付能力来自流程标准化和工程数据闭环。单个项目交付成功并不足以构成长期壁垒，持续承接多柜、多客户、多机房项目，需要厂商把 BOM、功率预算、机柜布局、烧机记录、故障日志、运输状态和现场部署问题沉淀为工程数据库。后续扩容时，服务器厂商可以基于既有数据优化设计规则、测试标准、包装运输和运维清单。超节点竞争由此从产品能力延伸到整柜测试、现场部署和项目复制能力。

图93：机架集成流程覆盖设计、装配、烧机、交付与现场部署

1. Rack Design	• Project inception & application analysis • Power budget • BOM creation • Rack layout
2. Quoting & Confirmation/Rack Prototyping	• BOM acceptance • Project initiation • Initial build • Proof of concept
3. Assembly	• Node assembly • Rack and stack • Networking and power • Cabling and labeling
4. Burn-In & Testing	• Multi-vendor interoperability • Full rack burn-in • Full rack test report • Performance report
5. Delivery	• Asset labeling and documentation • Crating • Air-ride trucking
6. On-Site Deployment	• SOW Documentation • OS/driver/software deployment • Deployment Testing
7. Signature Support	• Service level agreement • On-call support

数据来源：Supermicro 官网，东吴证券研究所

9. 风险提示

算力需求不及预期风险：若 AI 技术迭代放缓、商业化落地不及预期或下游资本开支收缩，将导致算力集群采购需求下降，影响超节点产业链订单与业绩兑现。

技术路线与商业化落地风险：超节点架构仍处快速演进阶段，核心技术路线尚未定型，若出现颠覆性方案或主流路线切换，或关键技术商业化验证周期超预期，将制约产业落地节奏，前期研发投入或面临过时风险。

整柜工程量产交付不及预期风险：超节点规模化交付需突破液冷、供电、系统验证等多重工程瓶颈，若厂商工程能力不足或供应链配套不及预期，将导致量产进度滞后、交付延期。

行业竞争加剧与盈利水平下滑风险：超节点赛道参与者增多，或引发价格竞争；叠加下游客户议价能力较强，行业盈利水平或承压，压缩产业链利润空间。

免责声明

东吴证券股份有限公司(以下简称“本公司”)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。

本研究报告仅供符合中国证监会《证券期货投资者适当性管理办法》规定的专业投资者使用,且接收人须为本公司认定并建立服务关系的专业机构投资者使用,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。本报告不面向普通投资者发布或传播。任何非专业机构投资者收到本报告后,应立即退回并销毁,不得依据本报告内容做出任何投资决策。

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>