

电子行业深度报告

超节点系列报告一：国产超节点方案量产元年，看好以太网成为主流技术路径

增持（维持）

2026 年 04 月 07 日

证券分析师 陈海进

执业证书：S0600525020001

chenhj@dwzq.com.cn

证券分析师 李雅文

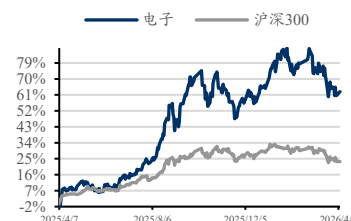
执业证书：S0600526010002

liy@dwzq.com.cn

投资要点

- **超节点：AI 算力基建核心形态，国内建设全面提速。**超节点作为整机柜级一体化紧耦合算力系统，已成为 AI 大模型训练与超算中心建设的核心基础设施。通过芯片级高速互联、统一液冷供电与集中式管理，超节点方案有效解决了传统服务器集群的带宽、时延与运维瓶颈。当前，全球 AI 算力架构加速向超节点切换，英伟达 GB200 NVL72 标杆方案落地的同时，国内华为 Atlas A3 900 SuperPoD、阿里磐久、中科曙光 ScaleX640 等厂商自研超节点产品密集推出，在集成密度、通信速率、计算能力等多项指标上接近国际一流水平。
- **Scale Up 主流协议：封闭 NVLink 与开放以太网双轨并行。**国际层面：英伟达 NVLink 封闭生态与 AMD/博通以太网开源开放双轨竞争。以 AMD 主导的 UALink、博通推出的 SUE 为代表基于以太网的开源协议，成为打破 NVLink 垄断的核心力量。**国内层面：**围绕自主可控、适配国产算力生态形成三条技术路线：自主可控专用系统总线（华为灵衢、海光 HSL）、以太网优化（字节跳动 EthLink、腾讯 Eth-X）与开放基础设施架构（中国移动 OISA）。
- **以太网：Scale Up 协议的未来主流，生态与技术双轮驱动。**以太网凭借开放生态、低成本、多厂商兼容的核心优势，正加速成为 AI 算力 Scale Up 的主流互联协议。（1）开放生态重构 Scale Up 竞争格局，以太网成破局 NVLink 垄断最优解。在 Scale Up 互联领域，英伟达 NVLink 凭借封闭生态长期构建技术壁垒，但随着 AI 集群向万卡级规模化演进，封闭垄断的弊端持续凸显。以太网依托全产业链开放生态，正成为打破垄断的核心路径。（2）技术策略补足延迟短板，以太网突破 Scale Up 性能瓶颈。以太网原生存在协议层过多、通信低效导致的延迟劣势，但在网计算和计算与通信已形成有效解决方案。
- **Switch 芯片国产替代玩家有哪些？**独立交换芯片厂商方面，国内玩家主要有盛科通信、数渡科技和澜起科技，整体格局为盛科通信领跑，25.6T 交换芯片已量产并导入主流设备商供应链；数渡科技和澜起科技聚焦 PCIe，形成多点突破。**大厂自研芯片方面，**海光、华为、中兴和新华三等服务器产业链厂商均有布局。其中，华为、中兴和新华三已推出自研交换芯片，并配套自家交换机及超节点方案；海光与华为则通过建设开放生态协议来抢占产业生态制高点。
- **投资建议：**重点推荐盛科通信、海光信息，建议关注中兴通讯、澜起科技、万通发展（数渡科技）等。
- **风险提示：**AI 应用进展不及预期，技术发展不及预期，市场竞争风险。

行业走势



相关研究

《端侧 AI 周跟踪：Google 发布 Gemma 4，模型能力跃迁催化终端硬件升级周期》

2026-04-06

《海外算力周跟踪：光互联 CPO&OCS 共振，PCB 扩产加码迎技术升级新周期》

2026-04-06

内容目录

1. 超节点：国内超节点能力跻身世界一流	4
2. 为什么更看好以太网成为 Scale Up 主流协议？	6
2.1. Scale Up 领域主流通通信协议发展	6
2.1.1. GPU-CPU 间直连协议：PCIe、CXL	6
2.1.2. GPU 多卡间直接通信主流协议：NVLink、以太网	6
2.2. 以太网有望成为 Scale Up 领域主流协议	8
2.2.1. 开放生态塑造竞争力，延迟短板可被技术策略补齐	8
2.2.2. 博通：覆盖多场景以太网产品矩阵，Tomahawk 系列交换机带宽持续跃升	10
2.3. PCIe 演进：博通与 AsteraLabs 的互联生态布局	10
2.4. 私有协议（英伟达 NVLink）的优劣势	12
3. 国产替代玩家梳理	13
3.1. 独立交换芯片厂商	13
3.2. 大厂自研交换芯片	15
4. 投资建议	18
5. 风险提示	18

图表目录

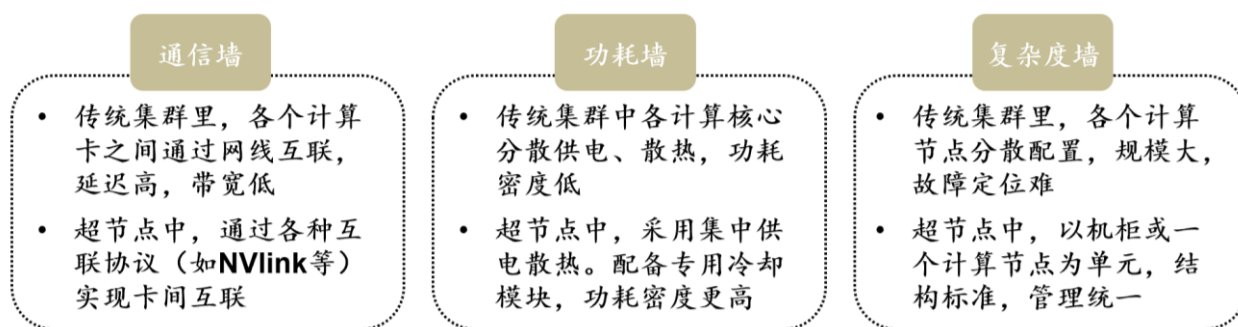
图 1: 传统集群面临的“三堵墙”问题.....	4
图 2: 超节点发展关键进程.....	5
图 3: PCIe、CXL 技术对比.....	6
图 4: NVLink、UALink、SUE 技术对比	7
图 5: UB 使用单一协议栈支持内存访问、消息传递、过程调用、资源管理	8
图 6: 博通 Tomahawk Ultra 使用在网计算技术大幅降低时延.....	9
图 7: 计算&通信重叠示意图	9
图 8: 博通以太网技术发展历程.....	10
图 9: PCIe 代际演进示意图.....	11
图 10: 博通与 AsteraLabs PCIe 相关业务线梳理	12
图 11: 盛科通信主要以太网交换芯片产品系列.....	13
图 12: 盛科通信 2018-2025 营业收入（亿元）	14
图 13: 数渡科技产品矩阵.....	14
图 14: 海光信息产品矩阵.....	16
图 15: 华为 2025 全联接大会产品路线图.....	17
图 16: 中兴通讯交换芯片与 NP 芯片技术路线图	17
图 17: 新华三智擎 660 网络处理芯片结构图.....	18

1. 超节点：国内超节点能力跻身世界一流

超节点是面向大模型训练与推理设计的新一代整机柜级一体化 AI 算力基础设施，它将数十到数百颗 GPU/NPU 在物理与逻辑层面深度紧耦合，通过自研高速互联协议、全局统一内存和整机柜级供配电与散热体系，让多机柜、上万颗芯片组成的集群在使用上如同一台巨型超级计算机，是当前高密度 AI 算力的主流架构形态。

引入超节点是为了解决传统服务器集群在大模型训练中遇到的**通信墙、功耗墙、复杂度墙**三大瓶颈：传统以太网与 PCIe 组网时延高、带宽不足，导致万卡集群算力利用率大幅下降；高密度 AI 芯片功耗飙升，风冷无法满足散热需求；松散堆叠的服务器部署复杂、故障点多、调度效率低，而超节点通过柜内紧耦合、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、简化部署，让万亿参数模型训练成为可能。

图1：传统集群面临的“三堵墙”问题



数据来源：华为 Atlas A3 900 SuperPoD 白皮书，东吴证券研究所

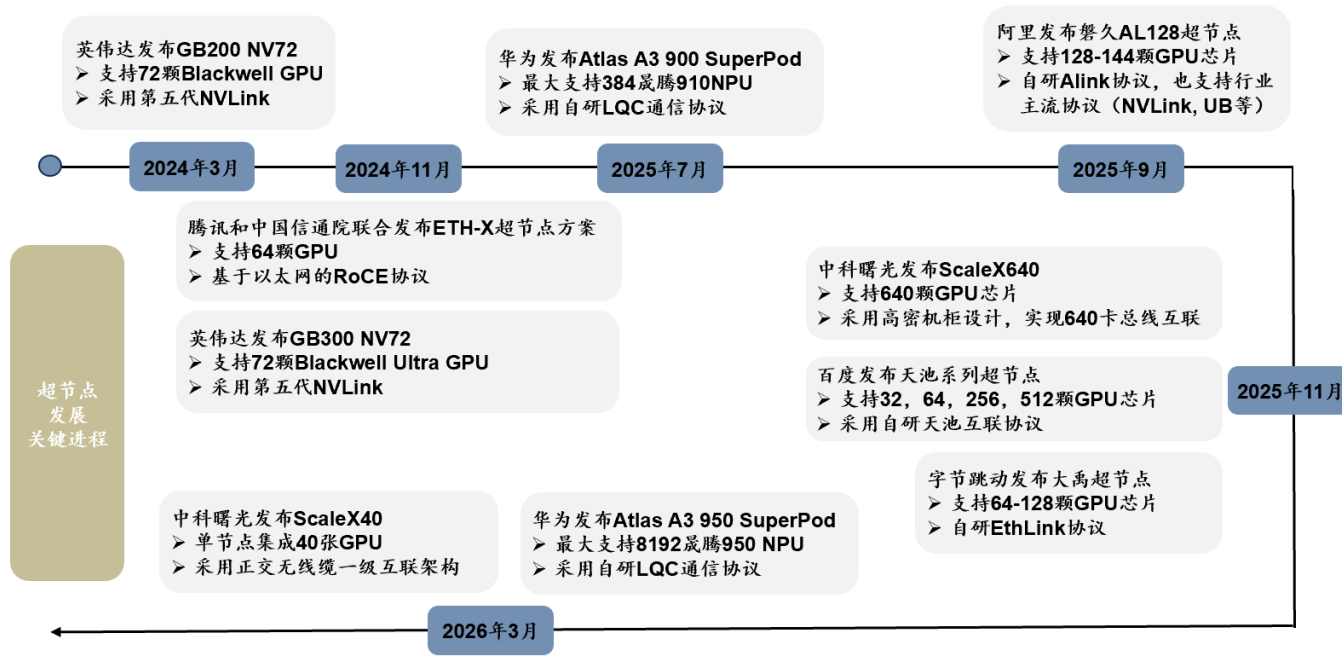
超节点的技术难点集中在高密集集成、高速互联、全局协同、可靠运维四个层面。首先，单机柜上百颗芯片带来极高的供电与散热压力，需要支撑数百千瓦级集中供电与高效液冷方案。其次，柜内 TB/s 级带宽、百纳秒级时延的紧耦合互联，对自研协议、交换芯片、高速线缆与信号完整性要求极高。同时，全局统一内存与硬件级集合通信需要软硬件深度协同；此外，高密度、高功耗、高复杂度的系统还必须解决漏液、供电冗余、故障定位与快速恢复等工程化难题，才能实现稳定商用。

海外层面，英伟达于 2024 年率先推出 GB200 NVL72 超节点，实现单机 72 卡高密度集成，依托 NVLink 高速互联协议实现节点内极低通信延迟与超高带宽。2025 年进一步迭代升级 GB300 系列超节点，持续提升算力密度与集群扩展能力。凭借集中式供电散热设计与全栈 CUDA 生态深度协同，成为全球超节点主流标杆方案。

国内层面，2024 年 11 月腾讯联合中国信通院发布 ETH-X 超节点方案，单机柜支持 64 颗 GPU，并通过以太网的 RoCE 协议实现机柜内互联；2025 年 7 月华为推出

Atlas A3 900 SuperPod，最大支持 384 颗昇腾 910 NPU 与自研 LQC 协议，9 月阿里发布磐久 AL128 超节点，以 128-144 卡高密配置、自研 ALink 协议实现 Pb/s 级带宽。2025 年 11 月，中科曙光 ScaleX640 以 640 卡规模刷新纪录，百度天池系列覆盖 32-512 卡多规格并搭载自研互联协议，字节跳动大禹超节点以 64-128 卡配置、自研 EthLink 协议实现标准化部署。2026 年 3 月，中科曙光推出 ScaleX40，华为推出 Atlas A3 950 SuperPod 等超节点方案，分别以 40 卡无线缆架构、8192 颗昇腾 950 NPU 的万卡级规模，推动国产超节点向更高算力、更大集群全面升级。

图2：超节点发展关键进程



数据来源：晟腾计算官网，英伟达 NV72 官网，阿里云博客，中科曙光官网，2025 中国智算中心全栈技术大会，中关村论坛年会，2025 百度世界大会，东吴证券研究所

2. 为什么更看好以太网成为 Scale Up 主流协议？

2.1. Scale Up 领域主流通信协议发展

2.1.1. GPU-CPU 间直连协议：PCIe、CXL

早期 CPU 与 GPU 互联以 PCIe 为通用基础、CXL 为内存一致性补充方案。二者依托成熟生态与技术演进，支撑了通用计算与中小规模多 GPU 系统的算力协同。两类协议虽满足基础互联需求，在大规模 AI 算力场景下存在性能与架构适配瓶颈。

PCIe (Peripheral Component Interconnect Express) 是通用型外设组件互连扩展标准。自英特尔于 2001 年提出，凭借长期迭代形成的成熟生态、良好的软硬件兼容性以及相对可控的部署成本，长期占据 CPU-GPU 及各类外设互联的主流地位，是服务器与加速卡互联的基础标配。但其树状拓扑结构决定了 GPU 之间无法直接进行高效 P2P 通信，跨 GPU 数据交互必须经过 CPU 或 PCIe 交换机中转，带来额外延迟与 CPU 资源占用。同时，PCIe 带宽与 GPU 本地 HBM 内存带宽存在量级差距，且协议本身不支持硬件级内存一致性，在多卡大规模扩展场景下效率明显受限。

CXL (Compute Express Link) 是基于 PCIe 的缓存一致性互联协议。通过新增缓存一致性与内存语义协议，实现 CPU 与 GPU 之间的内存共享与池化扩展能力，有效弥补了 PCIe 在异构内存协同方面的短板，为内存扩容与资源解耦提供了可行路径。但 CXL 整体性能仍受 PCIe 物理层带宽与延迟约束，实际访问效率存在明确上限；其核心优化方向集中在 CPU-GPU 内存访问，对 GPU 之间直接数据传输的提升作用有限，同时生态成熟度与大规模部署能力仍处于持续完善阶段。

图3：PCIe、CXL 技术对比

特性	PCIe	CXL
技术定位	通用外设/设备互连	新兴计算-内存一致性互连
语义	I/O 语义	内存语义
带宽	高	高
延迟	数百 ns	数百 ns
协议开销	低	中
可扩展性	有限	高
拓扑灵活性	中等	高
生态系统成熟度	高	发展中
典型使用场景	小规模多 GPU 系统、通用计算	解耦内存架构、GPU 内存扩展与共享

数据来源：宋晓勇等《扩展网络中节点内 GPU 互联调查：挑战、现状、见解及未来方向》，东吴证券研究所

2.1.2. GPU 多卡间直接通信主流协议：NVLink、以太网

国际层面：英伟达 NVLink 封闭生态与 AMD/博通以太网开源开放双轨竞争。

(1) 英伟达依托 NVLink 构建技术壁垒。新一代 NVLink5 交换机实现 14.4TB/s 无阻塞交换容量，单 Pod 可支持最多 576 个 GPU 互联；2025 年 5 月推出的 NVLink 融合

开放方案进一步将单端口带宽提升至 900GB/s，尝试兼容第三方定制 CPU/GPU，但底层技术的封闭性仍限制了多厂商 GPU 的深度协同。

(2) 以 AMD 主导的 UALink、博通推出的 SUE 为代表基于以太网的开源协议，成为打破 NVLink 垄断的核心力量。UALink 基于 Infinity Fabric 技术架构，2025 年 4 月发布的 1.0 版本单通道速率达 200GT/s，单 Pod 可支持 1024 个 GPU 互联，原生支持内存/显存共享与 Switch 组网，从架构层面实现多厂商 GPU 的协同适配，主打开源开放、多厂商兼容；博通 SUE 则深度绑定以太网技术生态，2025 年 6 月更新的规范支持 1/2/4 端口灵活配置，可适配全交换、网状拓扑等大规模组网场景，通过将 GPU 专有事务封装在以太网帧内复用现有协议逻辑，大幅降低部署成本，其与 UALink 协议层高度相似，但 SUE 更侧重现有以太网基础设施的兼容复用，部署门槛更低，适合大规模组网落地。

图4: NVLink、UALink、SUE 技术对比

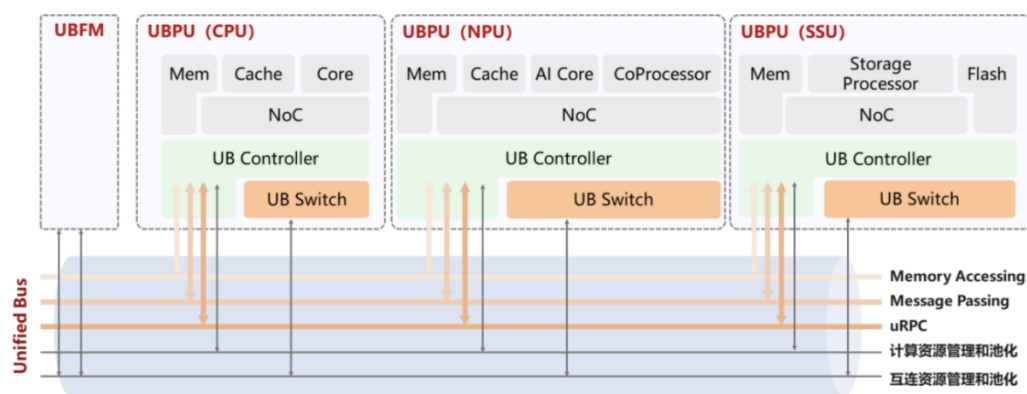
技术指标	NVLINK 5	UALINK	SUE
物理层	自研高速SerDes	基于以太网SerDes	完全兼容标准以太网SerDes
数据链路层	封装为小尺寸帧(含帧头、有效载荷、CRC校验)，支持链路层重试(Link Layer Retry,LLR)	封装64 B事务为640 B数据帧，添加CRC校验	保留以太网MAC层，新增10 B AI转发包头
传输层	事务层(支持内存读写、原子操作、消息传递等GPU协作所需的操作)协议层(支持NVSHMEM分布式共享内存)	事务层(压缩寻址，直接内存操作)和协议层(内存语义逻辑)	简化协议栈，直接处理内存语义通信，避免IP/UDP层开销
单向带宽	900 GB/s	640 GB/s	单通道200 Gbit/s，四通道800 Gbit/s
链路数量	每GPU 18条链路	x1, x2, x4	x1, x2, x4
互连规模	单Pod支持最多576个加速器	单Pod支持最多1024个加速器	单Pod支持512个加速器，扩展至1024个加速器
显存共享支持	完整支持，实现GPU间显存直接访问	完整支持，实现GPU间显存直接访问	不直接支持，通过集体操作卸载优化通信效率
生态兼容性	需专用交换机，不兼容PCIe/CXL等服务器协议	需专用交换机，兼容PCIe/CXL等服务器协议	完全兼容标准以太网设备和工具链

数据来源：张乾等《超节点关键技术与产业发展态势研究》，东吴证券研究所

国内层面：围绕自主可控、适配国产算力生态形成三条技术路线：自主可控专用系统总线（华为灵衢、海光 HSL）、以太网优化与开放基础设施架构。

(1) 华为灵衢（UnifiedBus, UB）：于 2019 年开始研究，随后发布灵衢 1.0 商用验证，于 2025 年 9 月发布并开放灵衢 2.0 技术规范。UB 协议栈由物理层、数据链路层、网络层、传输层、事务层、功能层以及 UMMU、UBFM 组成。单 Lane 速率达 118Gbps，支持 TB/s 级带宽互联；物理层线性直驱光技术实现<10ns 端口时延，链路层 Flit 传输优化时延与转发效率，整体达百 ns~us 级低时延。平等协同实现超节点全组件互联，无需 CPU 中转；支持多级 UB Switch 扩展及 UBoE 与以太网融合组网，适配超大规模集群扩张需求。

图5：UB 使用单一协议栈支持内存访问、消息传递、过程调用、资源管理



数据来源：华为《基于灵衢的超节点参考架构白皮书》，东吴证券研究所

(2) 海光 HSL (High-performance Scalable Link)：于 2025 年 9 月推出，12 月发布 1.0 规范。HSL 协议涵盖完整总线协议栈、IP 参考设计及指令集，海光已联合 10 余家厂商共建生态：全层级互联覆盖芯片、板卡、服务器等任意层级，支持全局地址空间与缓存一致性，降低编程复杂度，助力 CPU 与 AI 芯片“紧耦合”互联，加速国产算力集群建设。当前起步速率 112G，后续将升级至 240G，带宽较 32GPCIe 提升 8 倍，时延较 PCIe 降低约一半。

(3) 以太网技术路线以字节跳动 EthLink、腾讯 Eth-X 为代表。Eth Link 具备完整拓扑感知能力，覆盖服务器内及跨服务器 GPU 互联场景，单协议支持 1-4 个以太网接口，可通过低时延以太网交换实现最大 1024 个 GPU 节点的 Scale Up 组网，主打全场景拓扑覆盖、大规模组网能力；Eth-X 则针对 GPU-Switch 互通、GPU-GPU 访存两类核心需求分层设计，GPU-GPU 侧支持直接拷贝、DMA、统一编址等语义，GPU-Switch 侧基于以太网交换融合 DMA、ETH、MEM 能力，满足多层级传输需求，主打分层适配、多场景传输优化。

(4) 开放基础设施架构路线即中国移动研究院研发的 OISA 协议。定义物理层、数据链路层、事务层标准，具备统一报文格式、多语义融合、多层次流控等特性，通过两套报文格式适配不同应用场景，实现性能与资源利用率的协同优化，但当前与 CXL 标准的兼容性仍有提升空间。

2.2. 以太网有望成为 Scale Up 领域主流协议

2.2.1. 开放生态塑造竞争力，延迟短板可被技术策略补齐

开放生态重构 Scale Up 竞争格局，以太网成破局 NVLink 垄断最优解。在 Scale Up 互联领域，英伟达 NVLink 凭借封闭生态长期构建技术壁垒，但随着 AI 集群向万卡级规模化演进，封闭垄断的弊端持续凸显。以太网依托全产业链开放生态，正成为打破垄断的核心路径。(1) 英伟达亦加入产业协同：自 2023 年超以太网联盟 (Ultra Ethernet Consortium, UEC) 成立，而后在 2025 年 OCP 全球峰会上，英伟达也加入与 AMD、英

特尔、博通等巨头的合作，联合成立 ESUN 联盟（Ethernet for Scale-Up Networking），推动以太网适配 Scale Up 需求，打破 NVLink 单一生态绑定；（2）生态兼容规避供应商锁定风险，降低部署成本：以太网兼容现有数据中心基础设施，适配多厂商 XPU 设备协同，规避供应商锁定风险；同时兼容性减少定制成本，比如博通推出的 SUE（Scale Up Ethernet）框架完全兼容现有以太网交换机、光纤铜线、运维平台等基础设施，无需采购定制化专用芯片与设备，大幅降低部署成本。

技术策略补足延迟短板，以太网突破 Scale Up 性能瓶颈。以太网原生存在协议层过多、通信低效导致的延迟劣势，但在网计算与计算与通信重叠两大技术策略，已形成有效解决方案。（1）在网计算（In-Network Computing）：通过将部分计算任务卸载至网络设备（交换机、智能网卡）内部执行，替代传统主机端处理逻辑，减少数据往返主机的开销。博通 Tomahawk Ultra 交换机引入 In-Network Collectives（INC）能力，减少 All-Reduce 等通信量，将博通 Tomahawk5/Tomahawk6 以太网交换机时延从 600ns/800ns 优化至 250ns 级别，接近博通 PEX89000 系列 PCIe5.0 交换机 115ns 的低时延水平，提升集群整体计算效率。（2）计算与通信重叠（Computation-Communication Overlap）：通过调度优化实现计算与通信任务并行执行，打破“计算-通信-计算”的串行流程，避免 CPU/GPU 空闲等待。可与在网计算形成互补，进一步优化时延表现。

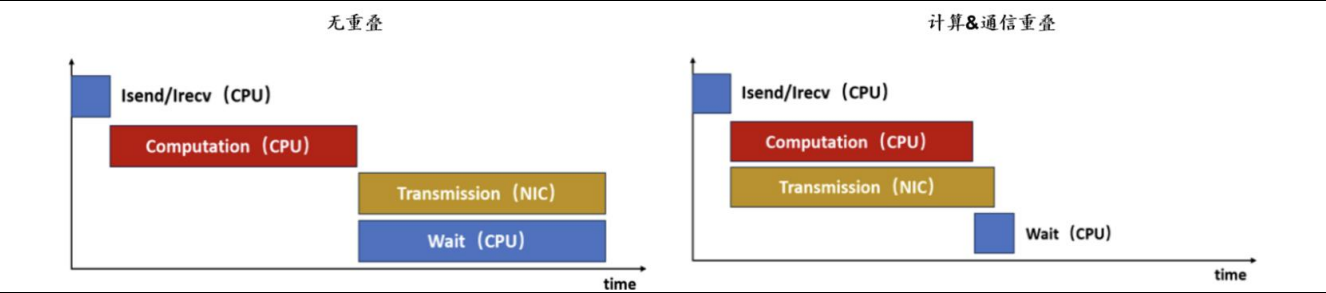
图6：博通 Tomahawk Ultra 使用在网计算技术大幅降低时延

博通TH5 TH6 TH Ultra（以太网）的时延 ≥ 250ns				博通PEX89000系列PCIe 5.0交换机时延 = 115ns					
Dimension	Tomahawk 5	Tomahawk 6	Tomahawk Ultra	Description	Lanes	Ports	Latency (ns)	Serial HPC	NT Ports
Primary Focus	High Throughput	Massive Scale	Ultra-Low Latency	PEX89144 Switch	144	72	115	✓	8
Process Node	5nm	3nm	5nm	PEX89104 Switch	104	52	115	✓	8
Bandwidth	51.2Tbps	102.4Tbps	51.2Tbps	PEX89088 Switch	88	44	115	✓	8
Latency	~600ns	~800ns	250ns	PEX89072 Switch	72	36	115	✓	8
Packet Rate	38 Bpps	76 Bpps	77 Bpps	PEX89048 Switch	48	48	115	✓	4
Ports Speed	800G/400G/200G	1.6T/800G/400G/200G	800G/400G/200G	PEX89032 Switch	32	32	115	✓	4
Lossless Mechanisms	End-to-end congestion control, adaptive routing, GLB	End-to-end congestion control, cognitive routing 2.0, GLB	LLR, CBFC, INC, Topology-aware routing	PEX89024 Switch	24	24	115	✓	4
Topology Support	Clos, torus, Dragonfly, Dragonfly+, Megafly	Clos, scale-up, rail-optimized, rail-only, torus	HPC-focused: Mesh, Torus, Dragonfly						
Use Case	Classic AI Scale-Out	Next-Gen Scale-Out & Scale-Up	HPC & AI Scale-Up						
Downward Compatibility	/	Tomahawk 5 software ecosystem compatible	100% pin-compatible with Tomahawk 5						

数据来源：FS 官网，博通官网，NETWORKWORLD，东吴证券研究所

注：右表通道数（Lanes）、端口数（Ports）单位均是个

图7：计算&通信重叠示意图



数据来源：Yuntian Zheng&Jianping Wu《Towards Efficient HPC: Exploring Overlap Strategies Using MPI Non-Blocking Communication》，东吴证券研究所

2.2.2. 博通：覆盖多场景以太网产品矩阵，Tomahawk 系列交换机带宽持续跃升

博通自 2007 年进军以太网领域，已构建覆盖多场景的以太网产品矩阵，并持续迭代核心产品推动技术向更高带宽演进。其以太网产线涵盖汽车 PHY/开关/微控制器、以太网网络适配器、以太网 PHY 和 PoE、以太网交换机及交换结构设备、以太网交换软件、以太网参考设计等，覆盖汽车、数据中心、企业等多元应用场景。Tomahawk 系列交换机带宽持续跃升，2016 年 Tomahawk2 达 6.4Tbps，2019 年 Tomahawk4 达 25.6Tbps，2022 年 Tomahawk5 推出，2025 年 Tomahawk6 达 102.4Tbps；同时 2024 年博通加入 UALink 联盟，2025 年推出 Scale Up Ethernet（SUE）框架规范，持续布局 AI 等新兴场景的以太网技术。

图8：博通以太网技术发展历程

年份	以太网相关里程碑事件
2007	展示并行光学以太网技术，具备100 Gbps的数据吞吐量
2008	首创推出下一代10 Gbps以太网SFP短距离收发器，用于网络设备
2010	• 制造了世界上第一颗Terabit级（Tbps）以太网交换芯片 • 发布市场首款四通道并行光学QSFP+ 能够实现多模式40 Gbps以太网上行的收发器应用
2012	引入了Vortex齿轮箱系列，包括28纳米CMOS 100 Gbps以太网/OTN物理接口
2016	首创搭载 Tomahawk 2（6.4Tbps） 以太网交换机的 64 个100GE 端口
2017	推出首款完全合规的TSN以太网交换机
2018	采样Thor，全球首个配备50G PAM-4和PCIe 4.0的200G以太网控制器
2019	出货了 Tomahawk 4 ，业界最高带宽的以太网交换芯片，速度达到 25.6 Tbps
2022	• 出货了Tomahawk 5 • 交付全球首个50G汽车以太网交换机
2024	• 交付了业界首个51.2 Tbps联合打包光学以太网交换机平台，面向可扩展AI系统 • 加入UALink联盟
2025	• 推出业界首款800G AI以太网网卡 • 宣布Tomahawk 6，业界首款102.4 Tbps以太网交换机，配备同封装光学器件 • 推出 Scale-Up Ethernet（SUE） 框架规范

数据来源：博通官网，UALink 官网，convergedigest，东吴证券研究所

2.3. PCIe 演进：博通与 Asteer Labs 的互联生态布局

PCIe 作为计算机系统核心 I/O 互联总线，历经多代技术迭代实现带宽与性能的持续跃升，成为支撑 AI 算力扩张的关键基础设施。从 2003 年 PCIe1.0（x16 双向带宽 8GB/s）起，PCIe 遵循每三年带宽翻倍的演进规律，通过编码方案、信号调制的持续优化实现性能突破：早期采用非归零调制（NRZ）与 8b/10b/128b/130b 编码，自 PCIe6.0 起引入 4 电平脉冲幅度调制（PAM4）与前向纠错（FEC），在相同频率下实现数据速率翻倍至 64GT/s，同时维持传输距离与 PCIe5.0 相当，PCIe6.0 x16 双向带宽达 256GB/s，后续 PCIe7.0/8.0 将进一步提升至 512/1024GB/s。尽管 PCIe 是 CPU-GPU 互联的主流标准，但其树状拓扑、主从（RC）架构为多 GPU 直带来带宽瓶颈，在 GPU 密集型扩展场景中，所有流量必须通过 CPU 或中央 PCIe 交换机，会带来不必要的延迟并浪费 CPU 资源。

图9: PCIe 代际演进示意图

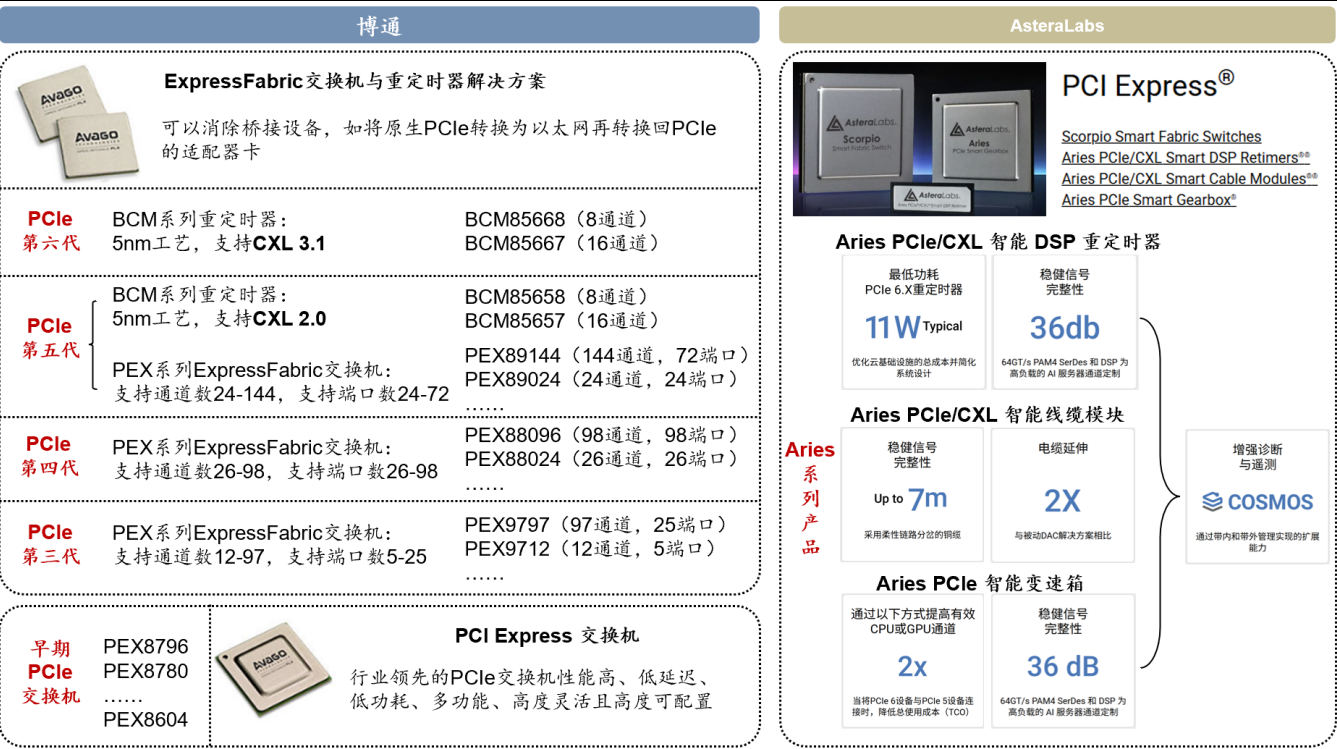
Version	Year	Signaling	Encoding	Max Data Rate Per Lane (GT/s)	Bidirectional Max Bandwidth (GB/s)				
					×1	×2	×4	×8	×16
1.0	2003	NRZ	8b/10b	2.5	0.5	1	2	4	8
2.0	2007			5.0	1	2	4	8	16
3.0	2010		128b/130b	8.0	2	4	8	16	32
4.0	2017			16.0	4	8	16	32	64
5.0	2019			32.0	8	16	32	64	128
6.0	2022	PAM-4 (FEC)	1b/1b (Flit Mode)	64.0	16	32	64	128	256
7.0	2025			128.0	32	64	128	256	512
8.0	2028 (Planned)			256.0	64	128	256	512	1024

数据来源: 宋晓勇等《扩展网络中节点内 GPU 互联调查: 挑战、现状、见解及未来方向》, 东吴证券研究所

博通: PCIe 行业的领导者, 已出货 10 亿端口, 围绕 PCIe 全代际构建全方位产品矩阵, 实现全场景覆盖。Gen6 层面推出 BCM85668、BCM85667 等 5nm 工艺 Retimer 芯片, 支持 PCIe6.0 与 CXL3.1 标准; Gen5 领域布局 BCM85657、BCM85658 等 Retimer 芯片, 以及 PEX89144、PEX89104 等多规格 Express Fabric 交换芯片, 端口数覆盖 24 至 144Lane, 适配不同规模数据中心架构; 同时向下覆盖 Gen4、Gen3 全系列以及早期 PCIe 产品, 如 PEX88064、PEX9716、PEX8796 等芯片, 满足存量市场升级需求。凭借行业领先的功耗控制、信号完整性与软件适配能力, 在 AI 行业, 博通产品广泛部署于全球超大规模运营商和系统 ODM/OEMs, 成为数据中心 PCIe 基础设施的核心供应商。

AsteraLabs: PCIe Retimer 行业的领先者, 聚焦 AI 与云基础设施核心场景, 以 PCIe6.0 连接产品组合为核心推进规模化部署。2019 年, AsteraLabs 采用新思科技的 DesignWare IP 推出 Aries Smart Retimer, 标志着业界首款 PCIe 5.0 Retimer 诞生。Aries 6 系列 Retimer 作为低功耗重定时方案, 典型功耗仅 11W, 具备 36dB 稳健信号完整性, 同时通过 COSMOS 诊断遥测技术提供全链路监控, 适配高要求 AI 服务器通道; 产品无缝兼容 Gen5 至 Gen6 代际架构, 可实现 3 倍传输距离扩展, 解决高速信号传输中的信号衰减、干扰问题。

图10：博通与 AsteraLabs PCIe 相关业务线梳理



数据来源：博通官网、AsteraLabs 官网，东吴证券研究所

2.4. 私有协议（英伟达 NVLink）的优劣势

NVLink 作为英伟达专属私有协议，实现了极致的性能与生态深度融合。它采用点对点直连架构，配合 NVSwitch 可构建非阻塞全互联拓扑，第五代 NVLink 单 GPU 双向带宽达 1.8TB/s，远超前代 PCIe 标准，同时借助 SHARP 技术在网络中完成数据聚合，大幅降低通信延迟；其支持内存一致性与统一内存池，让多 GPU 显存可被视作连续资源池，简化编程模型；且深度绑定 CUDA 生态，经长期调优后能充分释放算力，适配万亿参数模型训练等核心场景。

NVLink 存在严重的生态封闭性与高成本问题。它仅适配英伟达自家 GPU 产品，跨厂商硬件兼容性极差，易导致用户被单一供应商锁定；专用硬件（如 NVSwitch）与配套线缆价格高昂，大幅提升大规模部署成本；此外，其闭源特性限制了第三方定制与优化空间，在超大规模集群扩展中，相比开放标准协议的灵活性也稍显不足。

3. 国产替代玩家梳理

3.1. 独立交换芯片厂商

盛科通信：国内以太网交换芯片市场的领先企业，12.8T/25.6T 交换芯片已进入客户推广阶段。公司成立于 2005 年，凭借近 20 年的研发经验，已经开发出一系列覆盖接入层到核心层的以太网交换产品。其交换芯片产品覆盖 100Gbps-25.6Tbps 的交换容量以及 100M 到 800G 的端口速率，在企业网络、运营商网络、数据中心网络和工业网络等领域得到了广泛应用。如 TsingMa.MX 系列芯片交换容量为 2.4Tbps，支持 400G 端口速率，具备智能网络可视化技术和确定性网络技术；GoldenGate 系列芯片交换容量为 1.2Tbps，支持 100G 端口速率，特有的“灵云”设计为网络虚拟化应用提供了极具竞争力的硬件加速方案。其 TsingMa 系列芯片在中低端市场具有较强优势；在高端产品上，即数据中心领域，TsingMa.MX 系列芯片及 GoldenGate 系列芯片均已导入国内主流网络设备商并实现规模量产。且 TsingMa.MX 系列芯片还供货于新华三、锐捷网络和迈普通信，切入了国内主流设备商供应链。根据公司 2025 年中报，公司面向大规模数据中心和云服务需求，交换容量为 12.8Tbps 及 25.6Tbps 的高端旗舰芯片在客户处进入市场推广和逐步应用阶段，该产品支持最大端口速率 800G，搭载增强安全互联、增强可视化和可编程等先进特性。

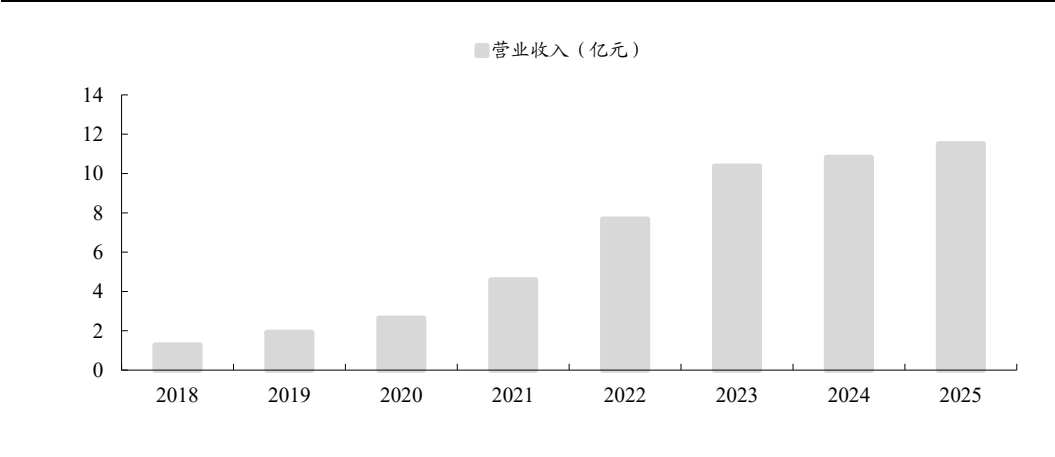
图 11：盛科通信主要以太网交换芯片产品系列

交换芯片	型号	图例	交换容量	最大端口速率	应用场景
-	-	-	25.6Tbps	800G	大规模数据中心和云服务
-	-	-	12.8Tbps	800G	大规模数据中心和云服务
TsingMa.MX	CTC8180		2.4Tbps	400G	企业网和边缘计算
TsingMa	CTC7132		440Gbps	100G	企业网、边缘计算和IP RAN
TsingMa.AX	CTC2118		28Gbps	1G	工业接入、SOHO/SMB 路由器接口扩展
TsingMa.CX	CTC5118		180Gbps	10G	企业网、边缘计算和工业网
GoldenGate	CTC8096		1.2Tbps	100G	架顶（ToR）/核心交换机 企业网汇聚/核心
GreatBelt	CTC5160		120Gbps	10G	城域网接入/汇聚 PTN/IP RAN 企业网接入/汇聚

数据来源：盛科通信官网，盛科通信公司公告，东吴证券研究所

盛科通信持续投入研发资源抢占国产替代空间。营业收入上，在 2019-2025 年间，营业收入高速增长，从 1.92 亿元上升至 11.51 亿元。同时以太网交换芯片业务在总营收中的占比也在不断提升。以太网交换芯片行业具备较高的技术壁垒、客户及应用壁垒和资金壁垒，因此当前行业整体国产程度低，国内参与厂商较少。盛科通信的以太网交换芯片在国内具备先发优势和市场引领地位，为我国数字化网络建设提供了坚实的芯片保障。

图12：盛科通信 2018-2025 营业收入（亿元）



数据来源：盛科通信官网，盛科通信公司年报，东吴证券研究所

数渡科技：主营业务为集成电路芯片设计，是国内极少数掌握 PCIe 5.0 交换芯片自主设计能力并可实现量产的企业。数渡科技 PCIe5.0 交换芯片具有高带宽、低延时、高可靠、兼容国际主流竞品、灵活可配置架构设计的特征，并支持互连芯片组网功能，兼容国内外主流的 CPU、GPU、DPU、SSD、网卡等芯片和设备，产品性能指标对标市场主流竞品，可实现对主流产品的兼容替代，且具有安全保密保障功能，填补了国内市场的空白。数渡科技产品支持片间组网功能，可实现 GPU 之间直接通信、协同工作，并支持构建大资源池和高可用集群，是国内构建自主超节点的稀缺选择，在国产替代领域形成先发优势，已经和业界头部客户厂家建立合作关系，样品已经通过主流厂家性能测试，目前已处于产品导入阶段，同时是国际 UEC 联盟、UALink 联盟和国内高通量以太网联盟成员。

图13：数渡科技产品矩阵

类别	产品名称	产品定位/说明
芯片	超节点智算服务器的高速交换芯片	突破国产智算服务器8卡限制，实现更大规模与更高性能
	通用服务器内部互连的PCIe/CXL高速扩展和交换芯片	连接ASIC、FPGA、存储设备等，实现高速互连与算力/数据处理提升
	存储服务器内部扩展的PCIe/CXL高速扩展和交换芯片	扩展连接高性能存储阵列，提升读写性能和存储密度
	边缘算力平台内部互连的PCIe高速扩展和交换芯片	连接传感器、控制器、存储器、网络设备，实现高效数据传输
智能算力系统方案	高性能算力平台系统方案	通过系统层面优化，在算力不足时提升整体效能
	高度智算服务器系统方案	突破单服务器资源限制，实现算力资源整合
	纵向扩展智算服务器模型	支持单台服务器纵向扩展至数十GPU/NPU
	高性能集中式存储服务器系统方案	支持PB级容量，RAID冗余，自动修复与安全功能
	全国异构边缘服务器系统方案	实现内存、计算、存储和网络资源池化与调度

数据来源：数渡科技公司公告，东吴证券研究所

澜起科技：在 PCIe 高速互连领域，澜起的产品版图不断拓展。以自研高速 SerDes 技术为支撑，继 PCIe 5.0 Retimer 芯片实现产业化之后，澜起科技新发布了 PCIe 6.x/CXL 3.x Retimer 芯片以及相应的 AEC 解决方案。同时，PCIe 7.0 Retimer 以及 PCIe Switch 芯片的研发正在有序推进，将为客户提供更全面、更具竞争力的 PCIe 互连芯片解决方案。2025 年澜起推出了符合 CXL 3.1 标准的 MXC 芯片，并已向主要客户送样；CXL 内存池化与扩展应用正在加速落地，AI 推理应用的到来将成为其规模部署的关键催化剂。目前，澜起正在研发以太网 Phy Retimer 芯片，基于在高速 SerDes 技术领域的长期积累与成熟的产业化经验，加快在该领域的产品布局。

云合智网：在以太网交换芯片领域，云合智网已形成 CLX84 和 CLX86 两大系列产品布局。CLX84 系列面向园区以太网交换机、企业数据中心和电信运营商汇聚交换设备，交换容量覆盖 2.4Tb/s 至 4.0Tb/s，集成高达 164 对 56G-PAM4 SerDes，支持 32×100GbE+4×200GbE、48×25GbE+8×100GbE 等典型接口配置，并提供灵活 I/O 和容量选择、业务可编程、端网安全、敏捷流量调度和智能运维等能力。CLX86 系列则面向大型数据中心、云计算和 AI 算力网络场景，交换容量提升至 8.0Tb/s 至 25.6Tb/s，集成多达 256 对 100G-PAM4 LR SerDes，支持 32×800GbE、64×400GbE、128×200GbE 以及 256×100GbE 等接口配置，强调单芯片高交换容量、高性能 SerDes、全共享缓存、灵活大表项、高效流量调度和智能运维等特性。两大系列产品均支持通用 SDK 软件开发套件和 OCP SAI。

楠菲微电子：在以太网交换芯片领域，楠菲微电子形成了企业网、城域网、工业以太网及高性能数据交换的多层次产品布局。SF2507 系列为支持 5+2 口千兆二层 SoC 以太网交换芯片，典型配置为 5 个 10/100/1000BASE-T/100BASE-FX 端口加 2 个 RGMII 接口；NF5120 面向企业网、城域以太网及工业场景，内置双核 600MHz 处理器，支持 8 路 SerDes 和最多 4 路万兆端口；NF5180 面向企业网、城域以太网和 SMB 场景，内置双核 A53 处理器，支持 12 路高速 SerDes 及最多 9 台设备堆叠；SF9056 系列定位工业以太网 SoC 交换芯片，集成 ARM 处理器，提供 56 路 SerDes 和最高 272Gbps 总交换容量，支持 L2/L3/L4 协议、隧道、SDN 可编程交换及 PCIe3.0×8 扩展；ES8000 系列则定位高性能、高连接性以太网交换芯片，单芯片集成 160 条 56G SerDes，提供最高 8.0Tb/s 带宽，支持 20×400GE、40×100GE/200GE 和 80×25GE/50GE/100GE 等典型配置，并支持 INT、OpenFlow 1.3、DCBX 和 FCoE 等特性。

3.2. 大厂自研交换芯片

海光信息：依托自研处理器与协处理器平台，在 Chiplet 互联和高速 I/O 方向深度布局，逐步构筑 Scale-Up 互联的核心能力。公司在 2025 年半年报中披露，公司在 CPU 和 DCU 芯片技术领域持续投入，已开展先进封装与高带宽低时延 Chiplet 互联研发，并通过 ComboPHY 支持处理器间互连总线、PCIe 及 CXL 等高速 I/O 接口，能够扩展片上网络（NoC）并强化 QoS 与低时延特性。2025 年 9 月 13 日，“海光系统互联总线协议开

放生态研讨会”在京举办。会上，海光面向 GPU、IO、OS、OEM 等产业全栈，正式宣布开放 CPU 互联总线协议（HSL）。据 25 年 12 月海光投关记录表介绍，公司与中科曙光保持深度战略协同，海光聚焦 CPU、DCU 芯片研发，曙光则在超节点算力、集群系统等前沿技术上深度合作，形成“芯片设计+系统集成”完整技术链条。海光 DCU 为 ScaleX640/ScaleX40 超节点提供自主可控算力核心，其 Chiplet 互联、ComboPHY 高速 I/O、HSL 等芯片级互联技术与曙光 ScaleFabric 高速网络形成从芯片内到系统间的全栈互联协同。我们认为公司同时布局 CPU、GPU（DCU）、Switch 互联，已初步形成算力基础设施全覆盖，随着 AI 大模型训练/推理需求增长，产品有望形成多点联动，卡位优势显著。

图14：海光信息产品矩阵



数据来源：海光信息公司公告，东吴证券研究所

华为：超节点集群布局领先，开源 Scale-out 互联协议抢占生态高地。华为于 1990 年开启交换芯片自研，1999 年自研 Solar 系列交换芯片，并于 2016 年发布 Solar 5.0 交换芯片。该芯片采用 16nm 制程，架构上的持续优化使之较上代版本提升了 4 倍的吞吐量。2024 年 9 月 19 日，华为在全联接大会上发布了单芯片 51.2T 数据中心盒式液冷交换机 CloudEngine XH9230。2025 年 9 月 18 日，华为举办全连接大会，发布新一代 AI 芯片、超节点与集群产品，同时开放“灵衢”互联协议。其中最新超节点产品 Atlas950 SuperPoD 和 Atlas 960 SuperPoD 分别支持 8192 及 15488 张卡，在卡规模、总算力、内存容量、互联带宽等关键指标上实现了全面领先。基于超节点，华为同时发布了超节点集群产品 Atlas950 SuperCluster 和 Atlas960 SuperCluster，算力规模分别超过 50 万卡和达到百万卡。会上，华为还宣布将开放面向超节点的互联协议灵衢（UnifiedBus）2.0 技术规范，邀请产业界伙伴基于灵衢研发相关产品和部件，共建灵衢开放生态。我们认为华为在 AI 算力基础设施从芯片、互联到整机集群整体布局，从产品路线图看已实现全栈打通，拥有持续翻倍迭代的能力。

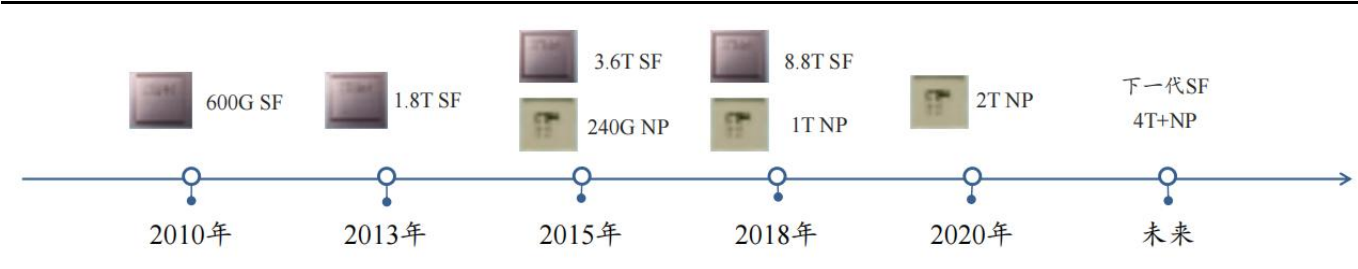
图 15： 华为 2025 全联接大会产品路线图

产品	上市时间	核心定位与亮点
Ascend 950PR / 950DT AI芯片	950PR: 26Q1; 950DT: 26Q4	新一代昇腾芯片，支持FP8/MXFP8/MXFP4/HiF8数据格式，互联带宽提升至2TB/s，分别面向推理Prefill与推荐场景（950PR）及推理Decode与训练场景（950DT）。
Ascend 960 / 970 AI芯片	960: 27Q4; 970: 28Q4	在算力、内存带宽、互联带宽等全面翻倍升级，960支持HiF4精度，970计划在FP4、FP8算力与互联带宽再翻倍。
Atlas 950超节点	26Q4	基于Ascend 950DT打造，支持8192卡，FP8算力8 EFLOPS，FP4算力16 EFLOPS，互联带宽16PB/s，训练性能比Atlas 900提升17倍。
Atlas 960超节点	27Q4	基于Ascend 960打造，支持15488卡，FP8算力30 EFLOPS、FP4算力60 EFLOPS，内存容量4460TB、互联带宽34PB/s，性能再度翻倍。
Kunpeng 950处理器	26Q1	96核/192线程与192核/384线程版本，支持通用计算超节点，安全新增四层隔离，首颗机密计算数据中心处理器。
TaiShan 950超节点	26Q1	全球首个通用计算超节点，基于Kunpeng 950，最大16节点、32颗处理器、48TB内存，支持内存/SSD/DPU池化，可平滑替代大型机和小型机。
灵衢（UnifiedBus）2.0互联协议	2025	面向万卡超节点的新型互联协议，实现高可靠全光互联、2.1微秒超低时延、TB级带宽，并首次向产业界开放2.0规范。
Atlas 950 SuperCluster集群	26Q4	64个Atlas 950超节点组成，50万+卡，FP8总算力524 EFLOPS，支持UBoE与RoCE协议。
Atlas 960 SuperCluster集群	27Q4	百万卡级集群，FP8总算力2 ZFLOPS，FP4总算力4 ZFLOPS，静态时延更低、可靠性更高

数据来源：华为全联接大会，东吴证券研究所

中兴通讯：已形成从通用高性能交换芯片到面向 AI 超节点的自研交换/NP 芯片全栈布局。在交换机芯片领域，中兴通讯在 2008 年启动交换芯片的自研，于 2011 年成功推出第一代自研交换网套片，并迅速在路由器等产品上成功应用。随后的几年，中兴通讯持续改进交换网技术，紧跟工艺革新的节奏，以 3 年一代的速度进行交换芯片的更新换代，陆续推出了 600Gbps、1.8Tbps、3.6Tbps 交换容量的 SF 系列交换芯片。并于 2018 年推出了 8.8Tbps 交换容量的第四代自研交换芯片。2020 年，中兴通讯启动了第五代自研交换芯片的研发。此外，中兴通讯还于 2015 年推出了首款自研 NP 芯片 SSP-1，并于 2019 年初推出了业界首款集成 FlexE 和 TSN 功能的 NP 芯片，展示了中兴通讯在交换机芯片技术上的领先地位和创新能力。2025 年，中兴通讯推出了基于自研 AI 交换芯片的超节点方案，GPU 间通信带宽达到 400GB/s 至 1.6TB/s，能够支持上百至上千张算力卡的高效互联。公司在超节点、高性能交换机、AI 交换芯片等关键环节完成了技术卡位，形成了国内市场的稀缺性。根据中兴通讯 2025 年年报，公司自研 DPU 定海芯片，支持 RDMA 标卡、智能网卡、DPU 卡等多种形态，提供高性能、多样化的算力加速硬件；携手产业伙伴推进 GPU 高速互联开放标准，打造正交超节点智算服务器，支持 Scale-Up 与 Scale-Out 双重扩展模式，具备业界领先的集成度与扩展能力。

图 16： 中兴通讯交换芯片与 NP 芯片技术路线图

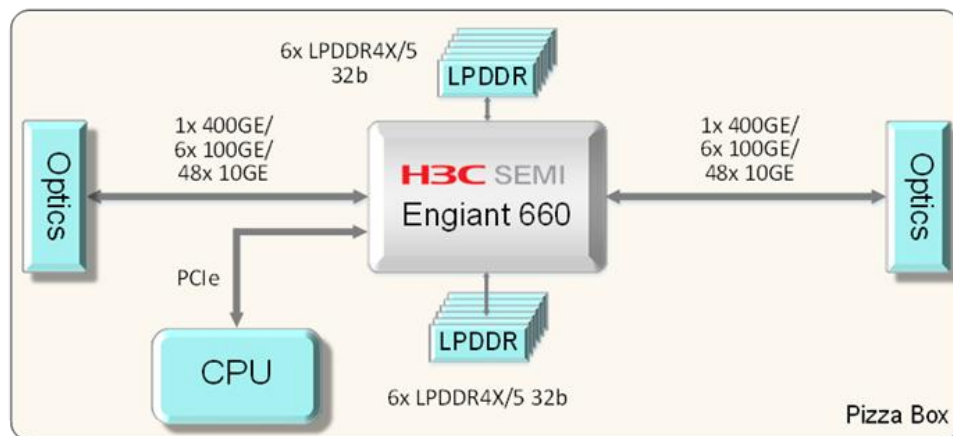


数据来源：中兴通讯官网，东吴证券研究所

新华三：自主研发智擎系列可编程 NP 芯片，接口吞吐能力 1.2Tbps。新华三于 2021 年 4 月正式发布其首款自研网络处理芯片智擎（Engiant）600/660，并于同年 7 月实现

量产。该系列芯片基于 16nm 工艺制造，集成 256 个专用 CTOP 处理核心、支持 4096 硬件线程，提供 1.2Tbps 的接口吞吐能力，集成度超过 180 亿晶体管，搭载 12 路 LPDDR5 控制器。芯片支持 C 语言编程，覆盖 L2-L7 层网络业务处理，内置流量管理、查找引擎、报文管理等加速单元，可灵活满足路由、交换、防火墙、负载均衡、SDN/NFV 等多类网络设备需求。智擎 660 有望推动国产高端网络处理器在运营商核心网、数据中心及智能安全设备领域的应用落地。在高效多元算力方面，H3C UniServer G7 系列服务器以丰富的软硬件生态兼容性，适用于集群训练、训推一体、边缘推理等场景。

图17：新华三智擎 660 网络处理芯片结构图



数据来源：半导体行业观察，东吴证券研究所

4. 投资建议

重点推荐盛科通信、海光信息，建议关注中兴通讯、澜起科技、万通发展（数渡科技）等。

5. 风险提示

AI 应用进展不及预期。若大模型相关应用在商业化落地、用户渗透或盈利模式验证上推进较慢，可能导致下游客户对算力基础设施的投入节奏放缓，从而影响超节点及相关产业链的需求释放。在此情况下，行业整体增长动能或阶段性承压，相关企业业绩兑现亦可能受到一定影响。

技术发展不及预期。超节点及高速互联等核心技术仍处于持续演进阶段，若关键技术在性能优化、稳定性或大规模部署能力上未能如期突破，或国产替代相关技术成熟度与生态完善程度不足，可能对产品落地及行业推广形成一定制约，进而影响产业化进程与市场空间释放。

市场竞争风险。随着行业关注度提升及参与者增加，国内外厂商在技术路线、产品性能及生态构建等方面的竞争或持续加剧，行业可能面临价格压力与份额分化风险。在竞争加剧背景下，部分企业盈利能力及市场地位可能受到影响，进而对行业整体回报水平产生不确定性影响。

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>