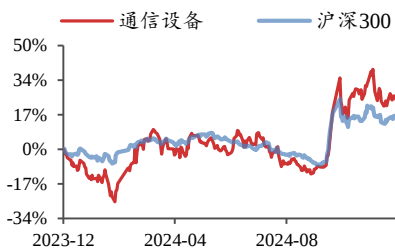


通信设备

2024 年 12 月 09 日

投资评级：看好（首次）

行业走势图



数据来源：聚源

相关研究报告

AI 时代交换机迎四大产业变革新机遇

——行业深度报告

蒋颖（分析师）

jiangying@kysec.cn

证书编号：S0790523120003

雷星宇（联系人）

leixingyu@kysec.cn

证书编号：S0790124040002

● 交换机变革 1：AI 驱动组网架构变革，新增后端组网需求

相比传统网络架构，AI 服务器组网增加后端网络组网（Back End），增加了每台服务器的网络端口数量，叠加 AI 集群加速 Scale out，万卡、十万、百万卡集群组网使得网络架构层数提升，有望带来大量高速交换机需求。当前 IB 仍然主导 AI 后端网络，但以太网根基深厚，生态厂商众多，未来以太网方案占比有望逐步提升，最终或将成为主流方案。

● 交换机变革 2：800G 交换机开始放量，102.4T 交换芯片有望推出

AI 大模型参数量持续增长倒逼集群规模提升，叠加 AI 芯片带宽提升，促使交换机端口速率及交换容量同步升级。交换机端口速率从 200G 向 400G、800G、1.6T 提升，交换芯片带宽容量提升至 25.6T、51.2T，下一代 102.4T 交换芯片有望于 2025 年下半年推出，盒式交换机端口数量得以持续增长以支持组网规模提升，高速数据中心交换机市场规模有望快速增长。

● 交换机变革 3：交换机白盒化趋势显著，带来新成长机遇

白盒交换机是一种硬件与软件解耦的网络交换机，其硬件由开放化的硬件组件组成，而软件可由用户或第三方自由选择和定制，具备灵活性、可扩展性较高、采购和维护成本较低等优势，广泛应用于互联网厂商和运营商网络，交换机白盒化趋势显著，目前产业生态较为完善，商用交换机芯片厂商、JDM/ODM/OEM 交换机设备商有望迎来发展新机遇。

● 交换机变革 4：光交换机商用逐渐成熟，光电融合组网落地大模型训练

光电路交换机（OCS）主要通过配置光交换矩阵，从而在任意输入/输出端口间建立光学路径以实现信号的交换，相比电交换机，光交换机具有成本低、时延低、功耗低、可靠性高等特点，在 AI 大模型预训练应用场景中表现较好。当前光电融合方案中 OCS 方案商用化程度较高，基于 3D-MEMS 系统的 OCS 方案综合应用较好。

投资建议：

AI 时代交换机迎来四大变革，交换机产业链有望长期受益。（1）交换机&交换芯片推荐标的：紫光股份、盛科通信、中兴通讯；受益标的：锐捷网络、菲菱科思、共进股份、烽火通信、Arista 网络、思科、Juniper、博通、Marvell 等；（2）全光交换机受益标的：光迅科技、Coherent 等；（3）工业交换机受益标的：映翰通、三旺通信、东土科技等；（4）交换机配套 AIDC 推荐标的：宝信软件、润泽科技；受益标的：光环新网、奥飞数据、云赛智联、网宿科技等；（5）交换机配套光器件推荐标的：中际旭创、新易盛、天孚通信；受益标的：华工科技、光迅科技、源杰科技等；（6）交换机配套液冷推荐标的：英维克；受益标的：科华数据、网宿科技、飞荣达、高澜股份、申菱环境等。

● 风险提示：云计算需求不及预期、数字经济增长不及预期、AI 发展不及预期

目 录

1、 交换机：AI 时代的核心网络通信设备	5
1.1、 交换机工作在数据链路或网络层，负责电/光信号转发.....	6
1.2、 交换芯片为核心部件，框式、盒式交换机各司其职.....	9
2、 交换机变革 1：AI 驱动组网架构变革，新增后端组网需求.....	13
2.1、 RDMA 技术被广泛应用于 AI 智算中心组网，IB 与以太网分庭抗礼.....	13
2.2、 AI/ML 后端市场快速增长，拉动交换机和网卡需求	20
2.3、 未来组网架构—Scale up 和 Scale out 的探讨	24
3、 交换机变革 2：800G 交换机开始放量，102.4T 交换芯片有望推出.....	26
3.1、 AI 大幅提升算力需求，驱动以太网交换机需求增长.....	26
3.2、 交换芯片不断升级，102.4T 芯片有望于 2025 年底推出	31
3.3、 AI 拉动高速交换机需求，全球 800G 交换机开始放量.....	33
4、 交换机变革 3：交换机白盒化趋势显著，带来新成长机遇.....	35
5、 交换机变革 4：光交换机逐渐成熟，光电融合组网落地大模型训练.....	38
6、 投资建议及交换机相关企业介绍	40
6.1、 盛科通信：稀缺的国产商用交换机芯片龙头	41
6.2、 紫光股份：国内交换机排头兵，率先布局 1.6T 端口、800G CPO 交换机.....	42
6.3、 中兴通讯：核心芯片自研，400G 交换机全场景布局.....	42
6.4、 锐捷网络：中标多个头部互联网客户项目，受益于白盒化浪潮.....	43
6.5、 共进股份：800G 交换机陆续交付，突破多个海外客户	44
6.6、 菲菱科思：发力中高端产品，国内领先 ODM/OEM 厂商	44
7、 风险提示	45

图表目录

图 1： 交换机种类较多	6
图 2： 交换机主要工作在 OSI 模型中的数据链路和网络层	7
图 3： 光信号通过端口进入交换机内部通过 PHY、MAC 层后进行转发.....	8
图 4： 美满 12.8T 芯片白盒交换机内部构成图	9
图 5： 信号在交换机中传递过程示意图	9
图 6： 以太网交换芯片报文处理架构	10
图 7： 典型以太网交换芯片架构	10
图 8： 博通 Trident 5 交换芯片架构.....	10
图 9： 框式交换机构造	11
图 10： 盒式交换机示意图	11
图 11： 框式交换网板示意图	11
图 12： 框式业务接口板示意图	11
图 13： 框式交换机逻辑结构示意图	11
图 14： 框式交换机正交架构信号传输轨迹.....	11
图 15： 园区、企业、工业场景网络组网架构.....	13
图 16： 智算中心网络以东西流量为主	14
图 17： 传统数据中心三层组网架构转向 Spine-Leaf 架构.....	14
图 18： AI 模型持续迭代	14

图 19: AI 模型训练所需算力持续增长	14
图 20: 网络带宽越高, 训练时间越少	15
图 21: 分布式全局加速取决于加速比	15
图 22: 集群规模与算力比率示意图	15
图 23: AI 集群组网可分为前端 (Front End) 和后端网络 (Back End)	16
图 24: 前后端网络组网均带来大量交换机需求	16
图 25: RDMA 通信时延更低	16
图 26: RDMA 协议栈	17
图 27: IB 在性能和功能上更优, RoCEv2 生态丰富具备成本优势	18
图 28: 以太网 RDMA 网络市场占比有望持续提升	18
图 29: HPC 集群中 IB 和以太网网络为主流方案	18
图 30: 多方助力 UEC 发展	19
图 31: GSE (全调度以太网网络) 分层架构	19
图 32: ETH+ 实现高载荷比、低时延的开放以太网	19
图 33: AI 服务器网络组网架构	20
图 34: RDMA 市场中交换机需求快速增长	20
图 35: AI 集群组网中计算和存储网络速率较高, 管理和业务网络速率较低	20
图 36: 主流三种智算网络架构适用不同场景	21
图 37: Spine-Leaf 两层架构组网	22
图 38: Fat-Tree 三层架构组网	22
图 39: DGX B200 网络端口示意图	22
图 40: DGX B200 127 节点计算网络组网架构	22
图 41: 算力集群拓展方向 Scale up + Scale out	24
图 42: Scale up——更多 GPU 互联	25
图 43: Scale out——更多节点互联	25
图 44: 数据中心以太网接口带宽向 800G、1.6T 发展	26
图 45: 工业自动化和汽车场景以太网速率较低, 运营商和云计算网络以太网速率需求较高	27
图 46: 交换机产业链涉及厂商众多	27
图 47: 芯片类在以太网交换设备的成本占比最大	28
图 48: 预计全球以太网交换设备市场 2020-2025 年复合增速为 3.2%	28
图 49: 预计国内以太网交换设备市场 2020-2025 年复合增速为 10.8%	28
图 50: 数据中心以太网交换机下游客户主要是云厂商和大型企业, 园区交换机客户主要为中大型企业	29
图 51: 2022 年国内数据中心交换机占比接近一半	29
图 52: 2021 年国内交换机制造以品牌商为主	29
图 53: 全球生成式 AI 数据中心以太网交换机市场规模有望快速增长	30
图 54: 2023 年全球以太网市场份额较为集中	30
图 55: 三种方式提升互联速率	31
图 56: 通道速率向 224Gbps 提升	31
图 57: 交换芯片迭代周期约为 2 年	31
图 58: 美满 Teralynx10 51.2T 芯片支持 32 个 1.6T 端口	32
图 59: 新华三首发 1.6T 端口智算交换机	32
图 60: 英伟达多芯片盒式交换机包含 72 个 1.6T OSFP 端口	32
图 61: 2023 年全球以太网交换机市场持续增长	33
图 62: 2024 年 800G 端口交换机有望加速放量	33
图 63: AI 后端网络速率有望加速迭代	33

图 64: 2023 年国内数据中心交换机市场持续增长.....	34
图 65: 2023Q1 国内新华三、华为、锐捷网络份额靠前.....	34
图 66: 白盒交换机在过去 30 年间加速发展.....	35
图 67: 白盒交换机硬件与软件解耦	36
图 68: 白盒交换机产业生态较为完善	36
图 69: SONiC 软件系统结构图	37
图 70: SONiC 参与厂商众多.....	37
图 71: 2021 年白盒交换机端口出货份额占比 10%	37
图 72: 2022 年白盒交换机端口出货份额占比 14%	37
图 73: 谷歌阿波罗架构核心层采用 OCS 交换机.....	39
图 74: Coherent 和光迅科技推出商用 OCS 产品	39
图 75: 交换机行业重要上市公司	40
图 76: 盛科通信产品包括交换芯片/模组、操作系统和交换机.....	41
图 77: 新华三数据中心交换机最大支持 51.2 万卡组网.....	42
图 78: 2023 年首发 800G CPO 交换机.....	42
图 79: 中兴 400G 框式及盒式交换机.....	43
图 80: 锐捷网络数据中心交换机产品矩阵.....	43
图 81: 锐捷网络推出 AI-Fabric 三级多轨网络架构.....	43
图 82: 共进股份推出数据中心交换机产品.....	44
图 83: 菲菱科思推出 12.8T 交换机.....	44
表 1: 以太网设备发展阶段	7
表 2: 交换机应用场景广泛	8
表 3: 框式、盒式交换机各司其职，满足不同场景需求.....	12
表 4: 框式交换机向零背板架构发展	12
表 5: 两层与三层无收敛组网的 GPU 最大节点数量.....	22
表 6: B200 组网服务器与交换机对应关系	23
表 7: WDM/MEMS 方案商用程度较好	38
表 8: 交换机行业受益标的估值表	41

1、交换机：AI 时代的核心网络通信设备

以太网交换机是重要的通信网络设备，随着全球 AI 的高速发展，AI 集群规模持续增长，AI 集群网络对组网架构、网络带宽、网络时延等方面提出更高要求，带动交换机朝着高速率、多端口、白盒化、光交换机等方向持续迭代升级，我们认为 AI 时代交换机有望迎来四大产业变革新机遇。

交换机变革 1：AI 集群新增后端组网需求，集群规模持续增长，以太网占比有望逐步提升，有望带来大量高速以太网交换机需求

(1) AI 训练集群带来 GPU 互联需求，新增后端网络组网需求。AI 服务器比传统服务器新增 GPU 模组，GPU 模组通过对应的网卡与其他服务器或交换机互联，实现各节点之间的通信。因此相比传统网络架构，AI 服务器组网增加后端网络组网 (Back End)，增加了每台服务器的网络端口数量，拉动对高速交换机、网卡、光模块、光纤光缆等组件需求。

(2) AI 集群加速 Scale out，万卡、十万、百万卡集群组网带来大量高速交换机需求。随着 AI 模型参数持续增长，带动集群规模从百卡、千卡拓展至万卡、十万卡，Scale out 推动组网架构从 2 层向 3 层、4 层架构拓展，带来大量高速交换机需求。

(3) 以太网网络根基深厚，生态厂商众多，AI 网络中以太网网络占比有望持续提升。IB 网络凭借低延迟、堵塞控制以及自适应路由等机制，仍然主导 AI 后端网络，但随着以太网网络部署的不断优化，超以太网联盟加速发展，我们认为未来以太网方案占比有望持续提升，带动以太网交换机需求增长。

交换机变革 2：AI 网络带来低时延、大带宽等网络需求，400G/800G 交换机持续放量，1.6T 交换机加速落地

AI 大模型参数量持续增长倒逼集群规模提升，叠加 AI 芯片带宽提升，促使交换机端口速率及交换容量同步升级。交换机端口速率从 200G 向 400G、800G、1.6T 提升，交换芯片带宽容量提升至 25.6T、51.2T，下一代 102.4T 交换芯片有望于 2025 年下半年推出，盒式交换机端口数量得以持续增长以支持组网规模提升，高速数据中心交换机市场规模有望快速增长。

交换机变革 3：交换机白盒化趋势显著，带来新成长机遇

白盒交换机是一种硬件与软件解耦的网络交换机，其硬件由开放化的硬件组件组成，而软件可由用户或第三方自由选择和定制，具备灵活性、可扩展性较高、采购和维护成本较低等优势，广泛应用于互联网厂商和运营商网络，交换机白盒化趋势显著，目前产业生态较为完善，商用交换机芯片厂商、JDM/ODM/OEM 交换机设备商有望迎来发展新机遇。

交换机变化 4：光交换机商用逐渐成熟，光电融合组网落地大模型训练

光电路交换机 (OCS) 主要通过配置光交换矩阵，从而在任意输入/输出端口间建立光学路径以实现信号的交换，相比电交换机，光交换机具有成本低、时延低、功耗低、可靠性高等特点，在 AI 大模型预训练应用场景中表现较好。当前光电融合方案中 OCS 方案商用化程度较高，基于 3D-MEMS 系统的 OCS 方案综合应用较好。

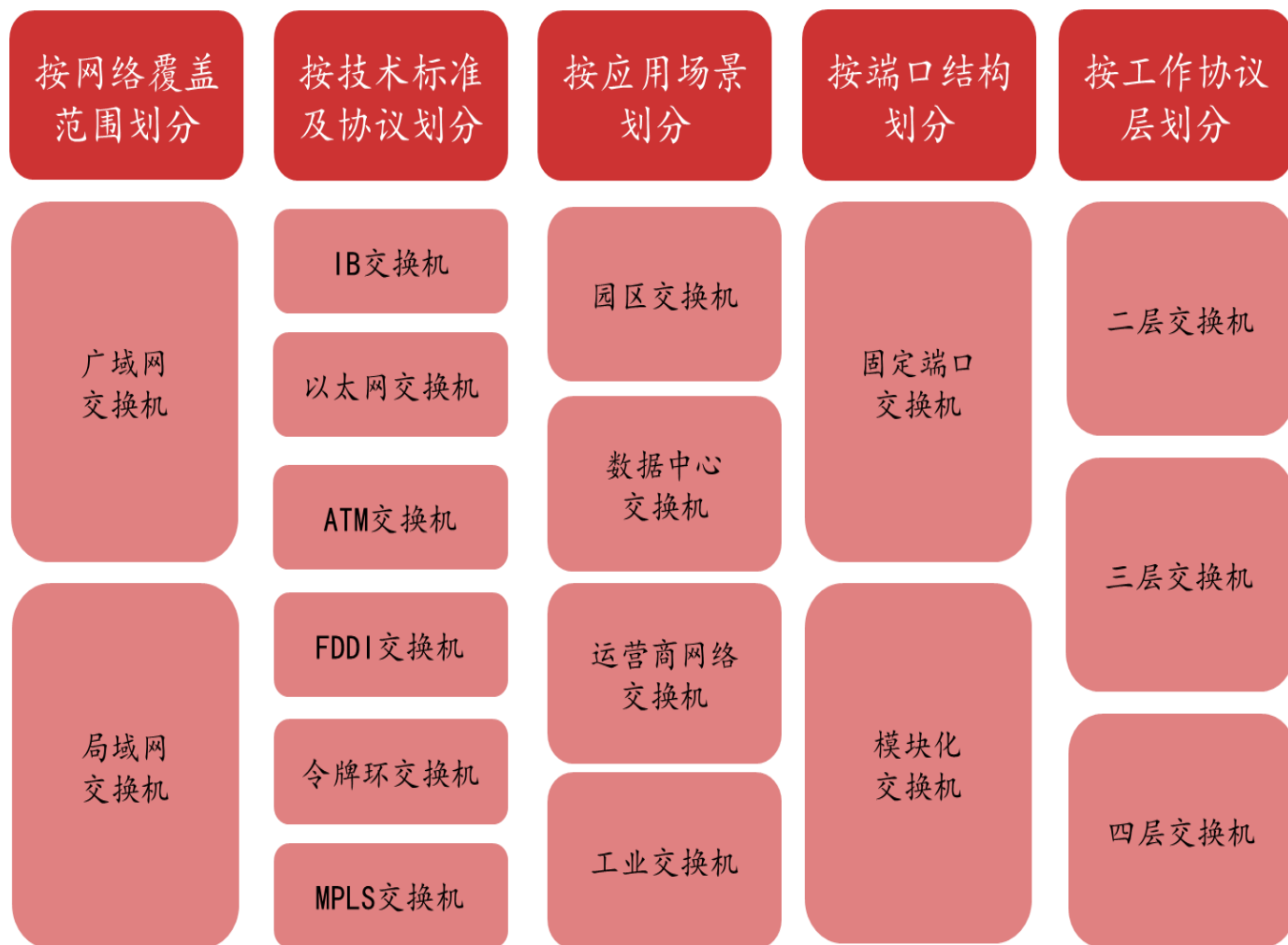
1.1、交换机工作在数据链路或网络层，负责电/光信号转发

交换机是用于电/光信号转发的网络设备。普通二层交换机（Switch）意为“开关”，是一种用于电(光)信号转发的网络设备。它基于 MAC 地址进行数据的转发，工作在 OS 七层模型中的第二层（数据链路层）。普通交换机具有多个端口，每个端口都具备桥接功能，可以连接一个局域网或一台高性能服务器或工作站。当设备接入交换机时，交换机会学习设备的 MAC 地址，并将 MAC 地址与端口对应起来，形成一张 MAC 地址表。在后续的数据传输过程中，交换机根据数据包中的 MAC 地址信息，将数据从对应的端口发送出去，实现数据的精准转发。

按应用场景和传输介质来看，交换机种类较多。交换机是重要的通信网络设备，最常见的网络交换设备以太网交换机为主，其次还包括语音交换机、光纤交换机等，适应不同网络环境与应用场景。

按照应用场景划分：（1）园区用以太网交换设备：可分为金融类、政企类、校园类；（2）运营商以太网交换设备：可分为城域网用、运营商承建用以及运营商内部管理网用；（3）数据中心以太网交换设备：可分为公有云用、私有云用、自建数据中心用；（4）工业以太网交换设备：可分为电力用、轨道交通用、市政交通用、能源用、工厂自动化用等。

图1：交换机种类较多



资料来源：头豹研究院、开源证券研究所

以太网交换设备已支持多个层级的数据转发，网络性能持续提升。早期以集线器为代表的以太网设备主要在物理层工作，无法隔绝冲突扩散，网络性能难以提升，而以太网交换机能够隔绝冲突，持续提升以太网性能。世界上第一台以太网交换机最早于 1989 年问世，经过三十余年的发展，以太网交换机在转发性能和功能上持续提升。转发性能方面，以太网交换设备的端口速率从 10M 发展到 800G，单台设备的交换容量从 Mbps 量级提升至 Tbps 量级。功能方面，以太网交换设备发展至今，可分为二层交换机、三层交换机和叠加型多业务交换设备。二层交换机和三层交换机之间的最大区别在于路由功能，叠加型多业务交换设备（四层或更高层）除了实现二层和三层的业务外，还可具备如防火墙、网关等其他功能。

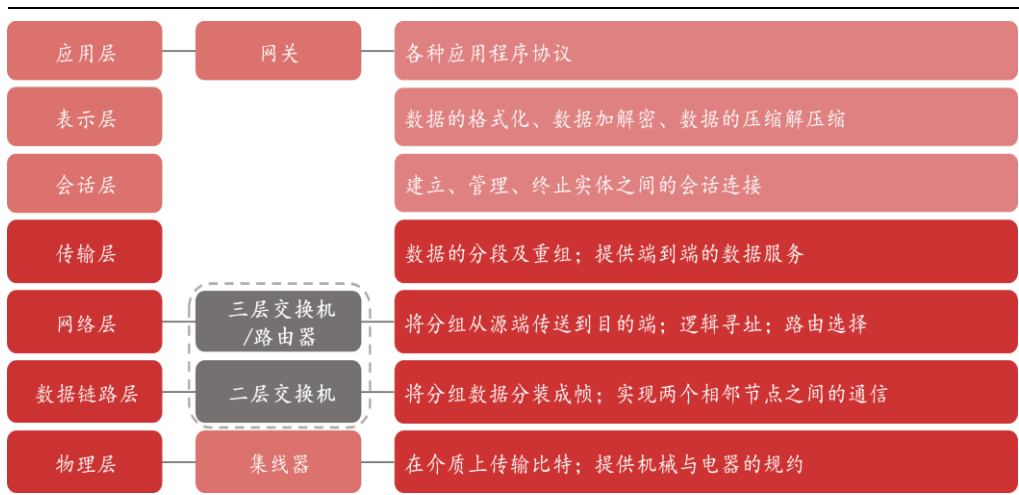
表1：以太网设备发展阶段

阶段	产品	转发硬件	应用场景
第一代	集线器	ASIC	共享式局域网
第二代	二层交换设备	ASIC	小型局域网
第三代	三层交换设备	ASIC	中小型局域网
第四代	叠加型多业务交换设备	ASIC+多核 CPU 混合模型	各类园区网、城域网

资料来源：灼识咨询、盛科通信招股书、开源证券研究所

以太网交换设备能够使不同网络中的设备终端实现互联互通。以太网交换设备对外提供高速网络连接端口，与主机和网络节点相连接，为接入设备的多个网络节点提供电信号通路和业务处理模型。以太网交换设备主要采用 OSI 模型，可作用于物理层、数据链路层、网络层、传输层或者应用层，通过高带宽的背部总线和内部交换矩阵实现多个端口对之间同一时间的数据传输和数据报文处理。

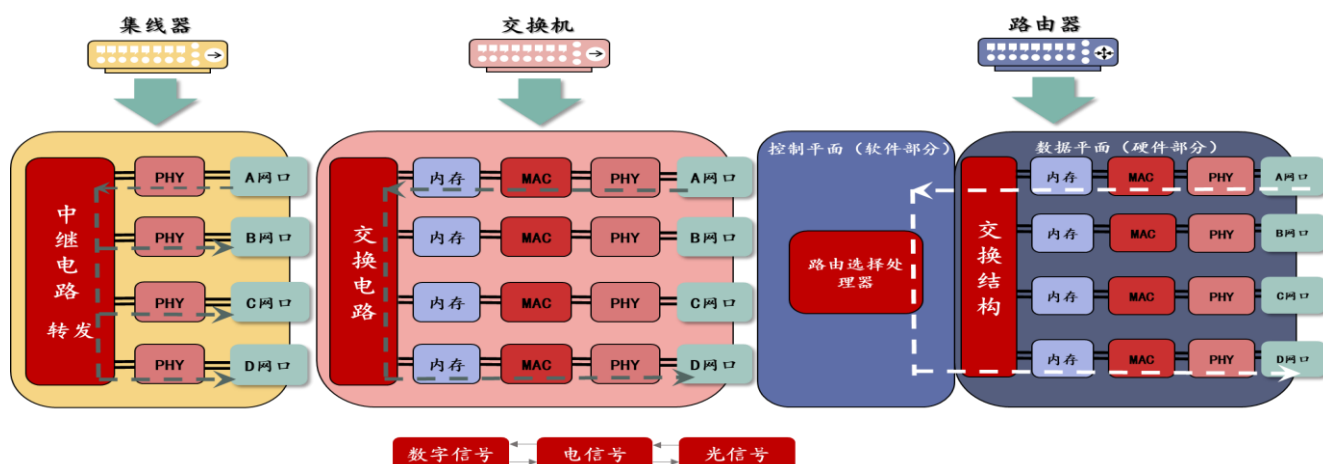
图2：交换机主要工作在 OSI 模型中的数据链路和网络层



资料来源：灼识咨询、盛科通信招股书、开源证券研究所

二层交换机工作在数据链路层，三层交换机工作在网络层。二层交换机在接受来自光纤传输的光信号后，通过光模块进行光电转换，最终将光信号转换为设备可理解的数字信号后，数据包从网络端口进入。PHY 层负责跨物理连接传输和接收比特流，包括编码、多路复用、同步、时钟恢复和线路上数据的序列化等，一旦在 PHY 上接收到有效的比特流，则数据将发送到 MAC 控制器，MAC 层负责将比特流转换为帧/数据包。经过以太网收发器芯片（PHY 芯片）、MAC 控制器后，进入以太网交换芯片，基于 MAC 地址进行数据交换；三层交换机/路由器工作在网络层，能够基于 IP 地址进行转发与路由选择。

图3：光信号通过端口进入交换机内部通过 PHY、MAC 层后进行转发



资料来源：开源证券研究所

交换机与其他网络设备功能各不相同。光猫：工作在物理层，通常安装在网络入口处即光纤接入点处，用于光电信号转换，主要应用在家庭网络接入场景；路由器：工作在网络层，连接不同网络，基于 IP 地址进行转发与路由选择；网关：通常用于连接使用不同协议的网络，能够在多个层次上进行必要的翻译和协议转换。二层交换机主要工作在数据链路层，具有网桥和集线器的功能，用于同一网络内基于 MAC 地址进行帧/数据包转发与过滤。

表2：交换机应用场景广泛

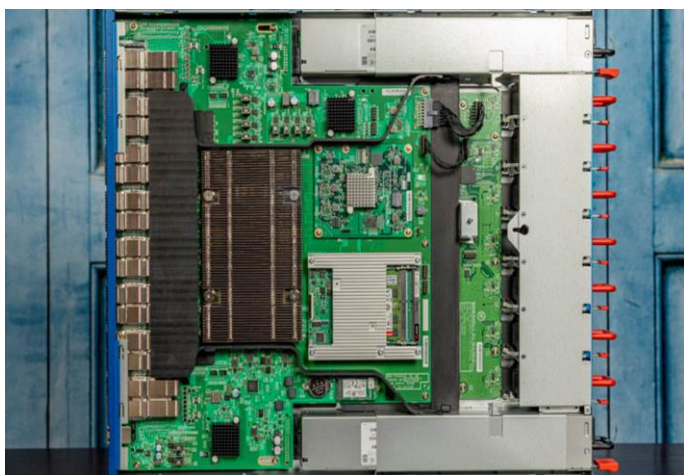
名称	OSI 模型层	主要用途	应用场景
光猫	物理层	作为光信号和电信号之间的转换桥梁	家庭宽带网络接入
网关	传输层/应用层	连接不同协议类型的网络	企业网、数据中心等
路由器	网络层	基于 IP 地址，实现数据包在不同网络中的转发	家庭网络、城域网、企业网、数据中心等
二层交换机	数据链路层	同一网络中基于 MAC 地址进行数据转发	家庭网络、城域网、企业网、数据中心等

资料来源：智慧傲联公众号、锐捷网络官网、开源证券研究所

1.2、交换芯片为核心部件，框式、盒式交换机各司其职

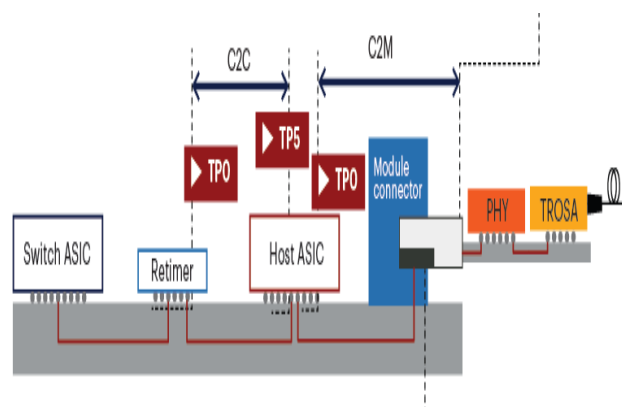
以太网交换机主要由芯片、PCB、光器件、插接件、阻容器件、壳体、电源、风扇等组成，芯片包含以太网交换芯片、CPU、PHY、CPLD/FPGA 等，其中以太网交换芯片和 CPU 是最核心部件。以太网交换芯片专为优化网络应用设计，是负责交换处理大量数据和转发报文的专用芯片，芯片内部的逻辑通路由数百个特性集合组成，以确保芯片在协同工作的同时保持较强的数据处理能力，架构实现较为复杂；CPU 是用于管理登录、协议交互的控制的通用芯片；PHY 负责处理物理层数据。

图4：美满 12.8T 芯片白盒交换机内部构成图



资料来源：STH

图5：信号在交换机中传递过程示意图



资料来源：Keysight

交换机的交换性能主要取决于背板带宽容量/包转发率、交换容量、端口速率和端口密度。背板带宽是衡量交换机数据吞吐能力的重要指标，其值越大说明该交换机在高负荷下数据交换的能力越强。在全双工工作模式下，当交换机的背板带宽容量 $\geq$ 交换容量（=端口数 $\times$ 端口速率 $\times 2$ ）时，才能实现线速转发（无阻塞转发），部分高端交换机采用无背板设计则需关注包转发率。一般来说，交换机拥有的端口速率越高则代表设备的处理性能越强，适用于数据流量大的场景；拥有的端口密度越大，则代表着设备的转发能力越强，可连接设备数量更多，组网规模更大。

以太网交换机芯片是以太网交换机中用于交换处理大量数据及报文转发的专用芯片，相当于网络方面的 ASIC，部分以太网交换机芯片内部会集成 MAC 控制器和 PHY 芯片。需要传输的数据包由物理端口进入以太网交换芯片后，芯片的解析器首先对数据包进行字段分析，为流分类做准备。通过安全检测的数据包进行二层交换或三层路由，流分类处理器对匹配的数据包作出相应动作，将可以转发的数据包根据 802.1P 或 DSCP 放到不同队列的 buffer 中，调度器根据优先级或 WRR 等算法进行队列调度并执行流分类修改动作，最后从端口发送该数据包。

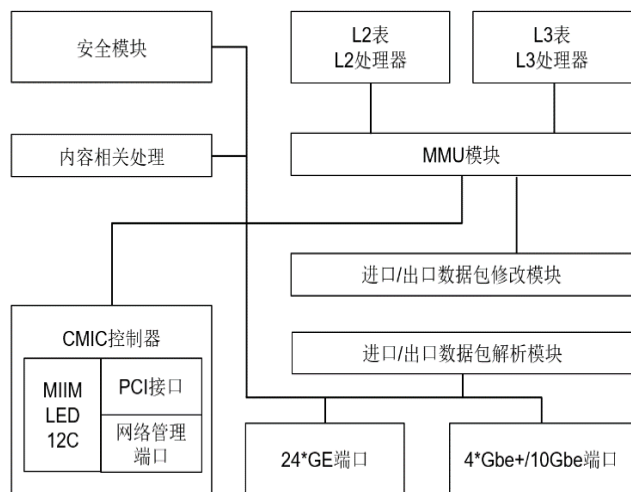
图6：以太网交换芯片报文处理架构



资料来源：灼识咨询、盛科通信招股书、开源证券研究所

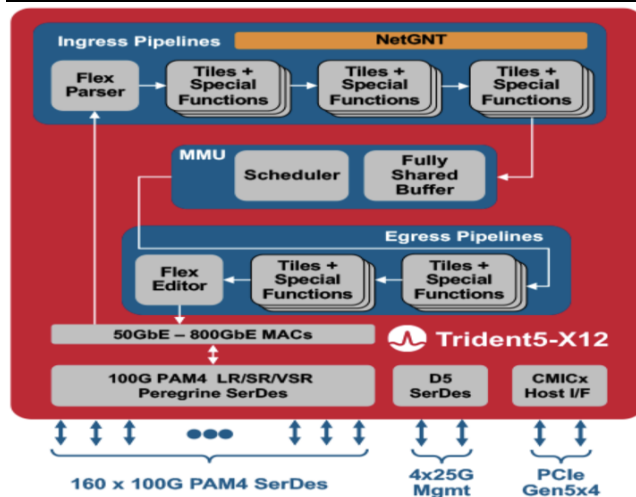
典型以太网交换芯片主要由接口模块、内容处理模块、进出口数据包修改模块、MMU 模块、L2 处理器（查阅 MAC 表）、L3 处理器（查阅路由表）、安全模块等模块组成，部分以太网交换机芯片内部会集成 CPU、MAC 控制器和 PHY 芯片。

图7：典型以太网交换芯片架构



资料来源：盛科通信招股书

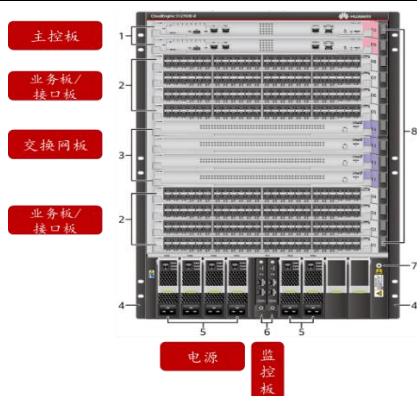
图8：博通 Trident 5 交换芯片架构



资料来源：博通官网

从交换机物理形态上，可以分为框式交换机和盒式交换机。框式交换机通常由一个机框和多个插槽组成，可以插入不同类型和数量的模块，如接口模块、主控模块、交换模块等，具有较高的灵活性和扩展性；而盒式交换机一般是一体化设计，接口数量和类型相对固定，部分盒式交换机接口采用模块化设计。框式交换机与盒式交换机的主要差异更多体现在内部构造与应用场景（OSI 使用层级）上。

图9：框式交换机构造



资料来源：华为官网、开源证券研究所

图10：盒式交换机示意图



资料来源：盛科通信官网

图11：框式交换网板示意图



资料来源：华为官网

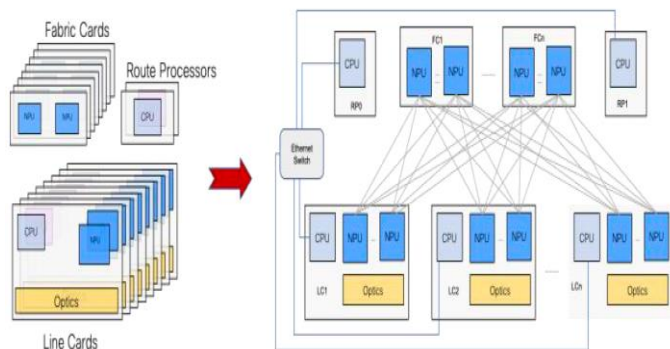
图12：框式业务接口板示意图



资料来源：华为官网

从应用场景来看，框式交换机背板带宽通常较大，可处理数据流量较大、适用端口密度更高的场景，并具备冗余机制，主要适用于大型企业/园区网络中的核心层、汇聚层，运营商网络核心、汇聚节点等，对性能和可靠性要求极高的场景；普通的盒式交换机，交换容量和端口速率有限主要用于中小企业/楼宇网络、边缘计算节点、中小型数据中心接入层；高速率盒式数据中心交换机，由于配备单个交换容量较高的交换芯片如博通 Tomahawk 5 或多个交换芯片如英伟达 Quantum-X800 Q3400 交换机系列，整体端口密度及端口速率较高，亦可用于中大型数据中心网络组网。

图13：框式交换机逻辑结构示意图



资料来源：ODCC《框式开放自研交换机技术实现与应用场景白皮书》

图14：框式交换机正交架构信号传输轨迹



资料来源：锐捷网络官网

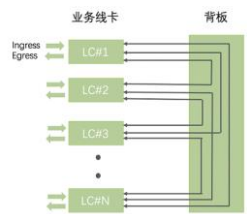
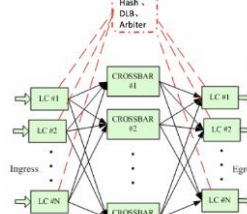
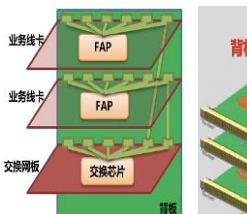
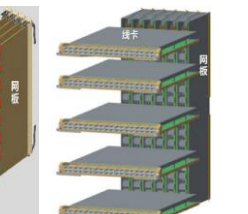
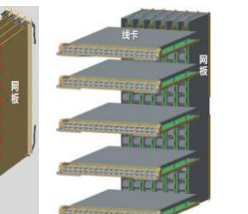
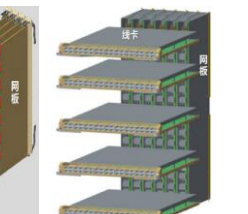
表3：框式、盒式交换机各司其职，满足不同场景需求

类别	框式交换机	盒式交换机
设计结构	模块化、多插槽	一体化设计、相对固定的接口数量
扩展性	高，灵活配置业务板卡	较低，灵活配置接口模块
性能与处理能力	强，适用于核心网络层	适中，适用于接入或汇聚层
电源模块	可配置，支持冗余	较固定，通常不支持冗余
堆叠技术	通常不支持	支持，可提升延展性但有限
升级能力	强，容易通过添加板卡升级	弱，通常难以升级
应用场景	大型数据中心、企业网络核心等	数据中心、中小型企业等

资料来源：锐捷网络官网、Networking Art、开源证券研究所

细分背板架构来看，常见的包括 Full-Mesh 交换架构、Crossbar 矩阵交换架构和 CLOS 交换架构等。(1) 在 Full mesh 架构中，所有业务线卡通过背板走线连接到其它线卡，任意两个节点间都直接连接，所以随着节点数量增加连接总数也持续上升，因此可扩展性较差；(2) Crossbar 架构则是一种两级架构，每个 CrossPoint 都是一个开关，交换机通过控制开关来完成输入到特定输出的转发，随着端口数量的增加，交叉点开关的数量呈几何级数增长，端口数量仍然有限；(3) CLOS 架构是一种多级架构，每个入口级开关连接至中间级开关再连接到出口级开关，每块业务线卡 and 所有交换网板相连，交换芯片集成在交换网板上，实现了交换网板和主控引擎硬件分离。CLOS 架构又可细分为非正交背板、正交背板和正交零背板设计。

表4：框式交换机向零背板架构发展

	Full-Mesh	Crossbar		CLOS		
架构图例						
分类	-	无缓存	有缓存	非正交背板	正交背板	正交零背板
硬件架构	无交换网板 线卡之间通过背板走线相连	单平面交换 集中调度 交叉点无缓存	单平面交换 分布式深度 交叉点有缓存	多平面交换 线卡和交换网板 平行 背板长走线	多平面交换 线卡和交换网 背板短走线	多平面交换 线卡和交换网板正交 无背板无走线
性能特点	受限于背板带宽和连接总数， 扩展性差 背板带宽是瓶颈	随端口数增加 CrossPoint 数量 呈几何增长 系统容量大时 仲裁器易形成 瓶颈	随端口数增加 CrossPoint 数量 呈几何增长 调度算法复杂 限制扩展	背板限制带宽扩展且无法实现直通散热 走线带来信号衰减 基于 cell 的动态负载实现无阻塞	背板限制带宽扩展且无法实现直通散热 基于 cell 的动态负载实现无阻塞	带宽扩展更换相应网板即可 无背板设计实现交换机直通散热 基于 cell 的动态负载实现无阻塞
适用设备	低密度槽位	高密度槽位 可面向未来 1-3 年扩展	高密度槽位 可面向未来 1-3 年扩展	高密度槽位 可面向未来 1-3 年扩展	高密度槽位 可面向未来 10 年扩展	高密度槽位 可面向未来 10 年扩展

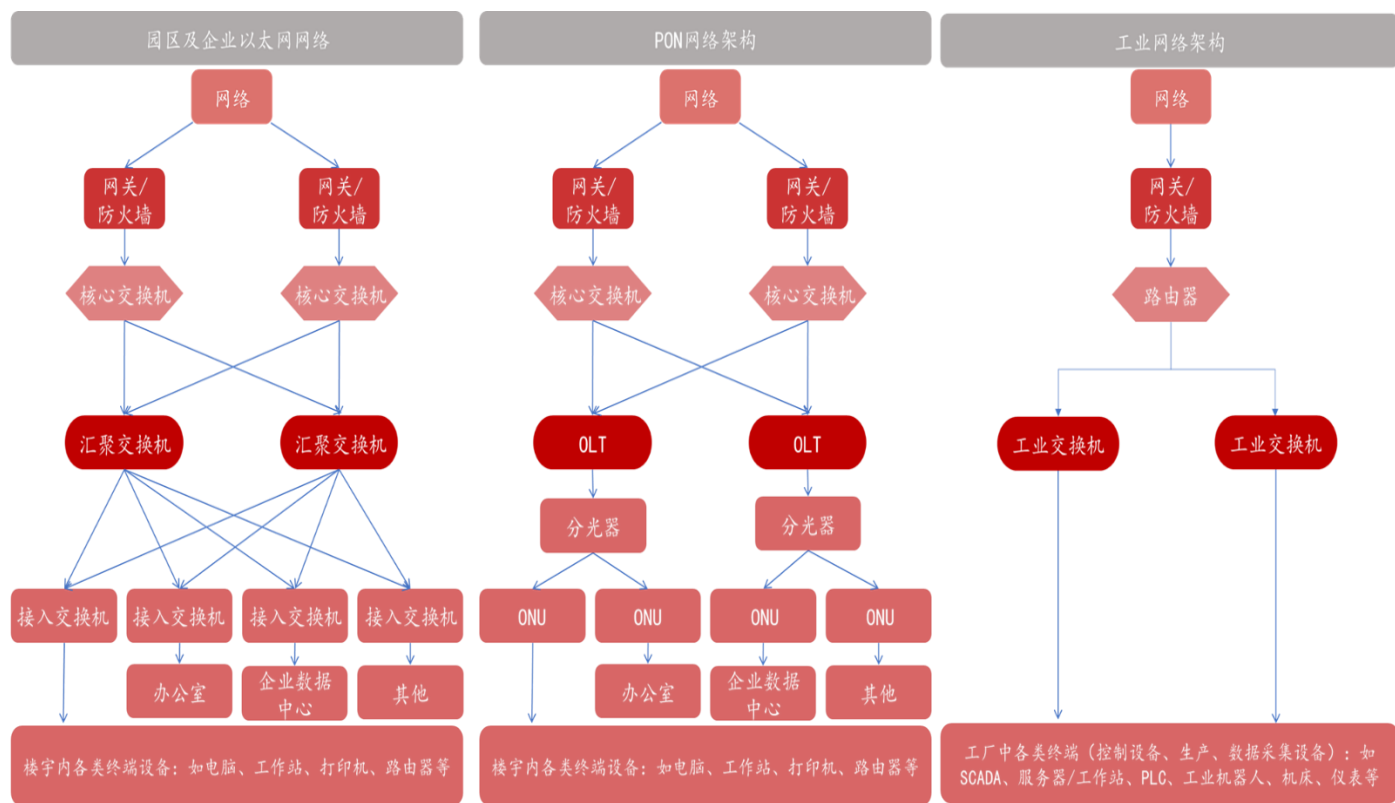
资料来源：锐捷网络官网、开源证券研究所

2、交换机变革 1：AI 驱动组网架构变革，新增后端组网需求

2.1、RDMA 技术被广泛应用于 AI 智算中心组网，IB 与以太网分庭抗礼

交换机下游需求场景主要包含数据中心、园区及企业、工业和运营商共 4 种网络场景，不同场景下交换机组网架构略有区别，根据终端设备数量采用 2 层或 3 层组网架构。其中，园区场景可采用无源光纤网络（PON）网络 2 层架构组网，主要由核心交换机、光线路终端 OLT、无源分光器 POS、ONU 光网络单元组成；也可采用全光以太网组网，由核心、汇聚、接入交换机组成。工业场景中对交换机速率要求不高，需要设备的稳定性和可靠性更高，能适应严峻的工业作业场景。

图15：园区、企业、工业场景网络组网架构

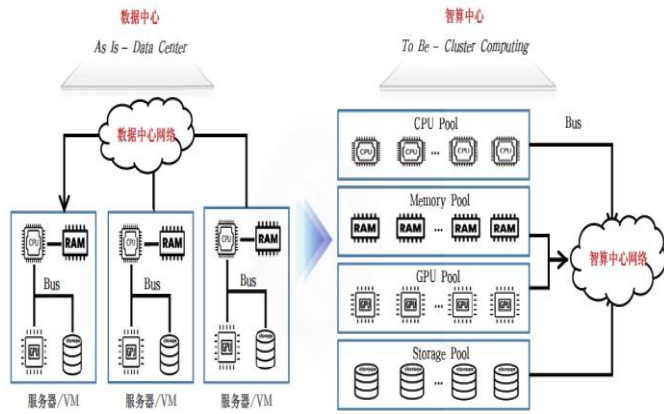


资料来源：锐捷网络官网、开源证券研究所

传统的数据中心主要面向业务场景，以服务器或虚拟机为池化对象，网络提供服务器或虚拟机之间的连接，数据大多进行南北向流动。而智算中心主要面向任务场景，以算力资源为池化对象，网络提供 CPU、GPU 和存储之间的高速连接，数据大多进行东西向流动。

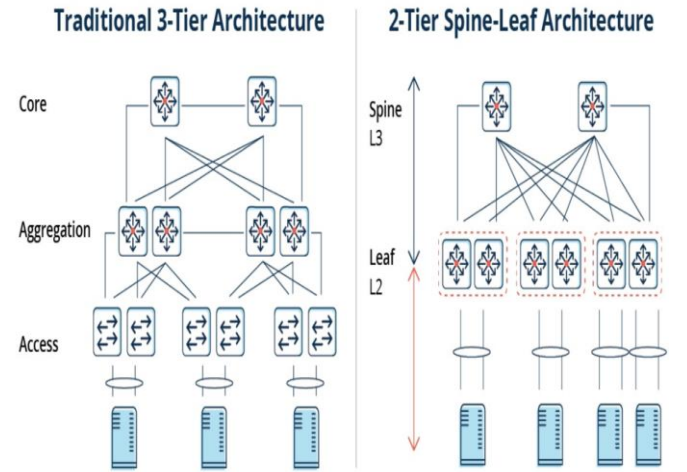
传统三层网络架构主要是为南北向流量设计，包括核心层、汇聚层、接入层交换机，当东西向流量较大时，服务器间通信数据转发路径有 5 跳，汇聚层和核心层交换机的流量压力会快速增加，网络性能会局限在汇聚层和核心层。相比 Spine-Leaf 架构，扁平化设计可缩短服务器之间的通信路径，转发路径仅有 3 跳从而降低延迟，更适合同样流量较大的业务，同时，Spine-Leaf 拓扑架构还拥有较好的可扩展性，横向扩展时不需要重新架构。

图16：智算中心网络以东西流量为主



资料来源：《智算中心 IB 及 RoCE 网络技术探究》李家清等

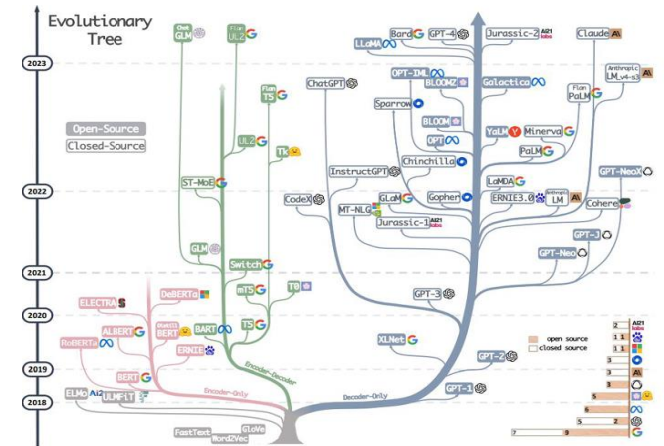
图17：传统数据中心三层组网架构转向 Spine-Leaf 架构



资料来源：HPE aruba 官网

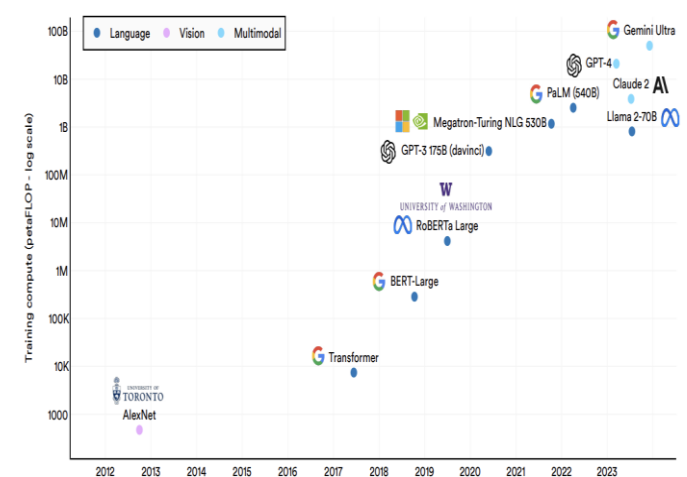
AIGC 发展带来蓬勃算力需求，大模型加速迭代拉动。2022 年底生成式 AI 大模型 ChatGPT 横空出世，掀起新的 AI 浪潮，海内外云计算厂商和研究院所等企业均陆续投入到大模型研发当中。人工智能发展迅速，AI 大模型快速迭代，从语言模型走向多模态，模型架构不断优化，出现 MOE 混合专家模型等架构。当前 Scaling Law 依旧成立，模型为获得更好的性能，数据量和参数规模均呈现“指数级”增长，算力需求持续增长。从参数量来看，以 GPT 模型为例，GPT-3 模型参数约为 1746 亿个，训练一次需要的总算力约为 3640 PF-days。据中国信通院数据，2023 年推出的 GPT-4 参数数量可能扩大到 1.8 万亿个，是 GPT-3 的 10 倍，训练算力需求上升到 GPT-3 的 68 倍，在 2.5 万个 A100 上需要训练 90-100 天。

图18：AI 模型持续迭代



资料来源：Harnessing the Power of LLMs in Practice: A Survey on ChatGPT and Beyond

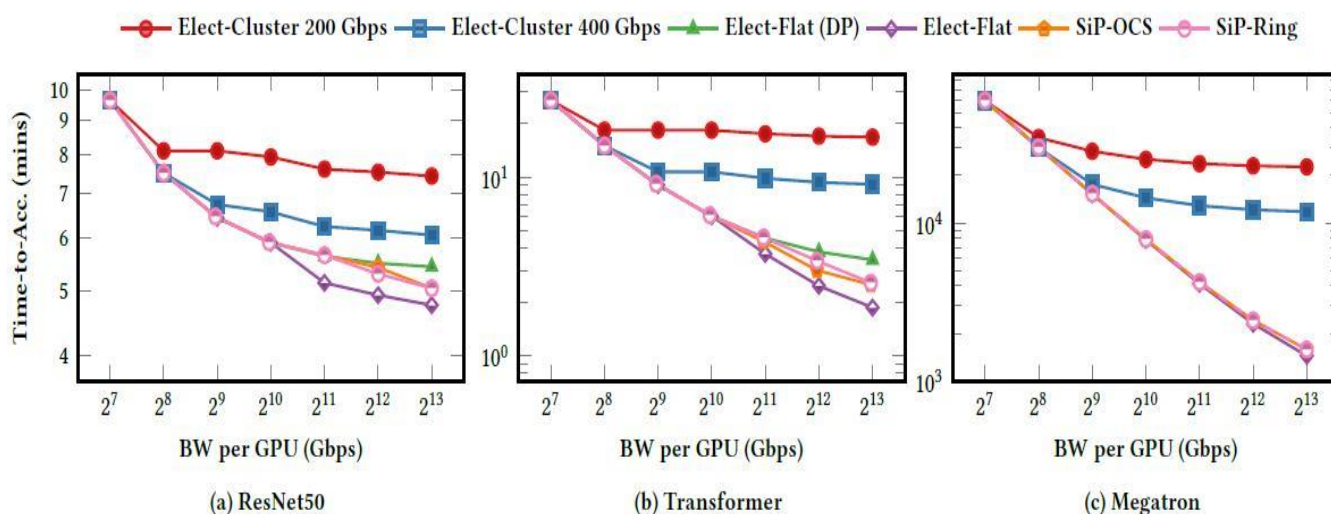
图19：AI 模型训练所需算力持续增长



资料来源：斯坦福 HAI 研究所

AI 模型参数持续增长，单卡算力和显存受限，迫使训练集群规模持续增长，AI 军备竞赛下，网络性能成制胜关键之一。对于参数较大的 AI 模型，更大的网络带宽能够显著减少训练完成时间。

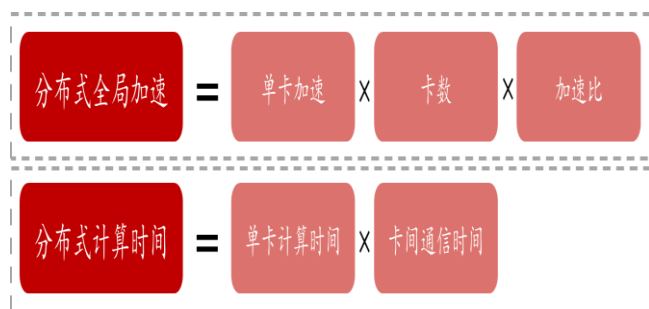
图20：网络带宽越高，训练时间越少



资料来源：ODCC《下一代以太网技术需求白皮书》

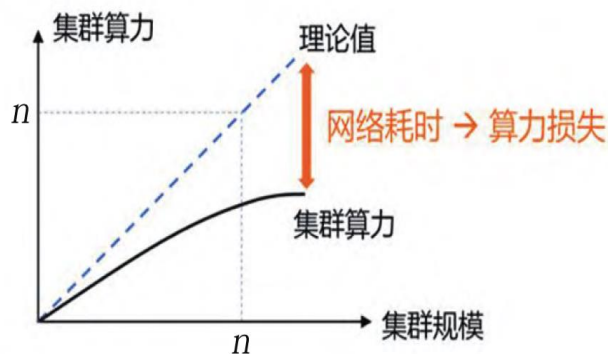
伴随模型参数增大，网络瓶颈问题逐渐凸显。AI 大模型以 GPU 集群分布式训练为基础，但随着 AI 算力集群规模的扩大，超大集群并不意味着超大算力，训练过程中各集群节点间需要频繁地进行参数同步，产生大量通信开销。**集群有效算力正比于（GPU 单卡算力×总卡数×加速比×有效运行时间）**，加速比指单处理器和并行处理器在处理同一任务时消耗时间的比率，当集群规模越大时，集群算力的实际增长程度越低，即集群算力不等于集群 GPU 数量与单 GPU 性能的乘积，集群算力的损失值大多是因为网络耗时导致，而网络性能则是决定 GPU 集群算力加速比的关键。

图21：分布式全局加速取决于加速比



资料来源：百度智能云《智算中心网络架构白皮书》、开源证券研究所

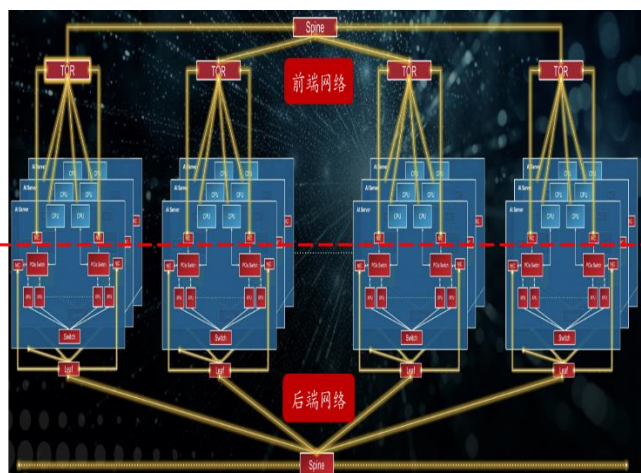
图22：集群规模与算力比率示意图



资料来源：李家清等《智算中心 IB 及 RoCE 网络技术探究》

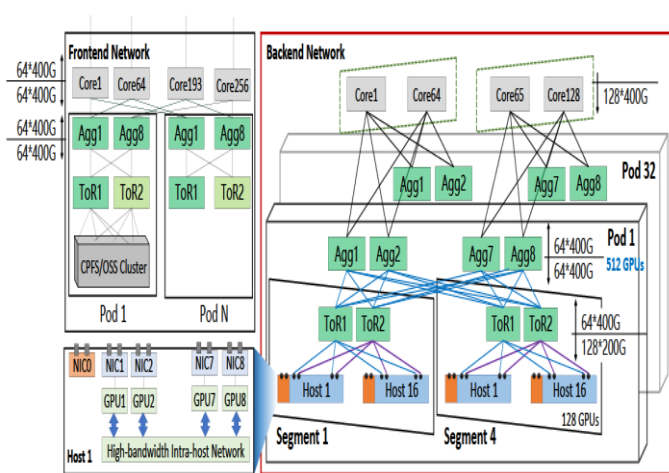
AI 模型参数持续增长，AI 训练集群带来 GPU 互联需求，新增后端网络组网需求。传统数据中心架构下，传统服务器与交换机之间通过网卡互相通信，网卡可直连 CPU 进行数据交换；AI 服务器比传统服务器新增 GPU 模组，服务器内部 GPU 之间通过 PCIe Switch 芯片或 NVSwitch 芯片实现内部互联，GPU 模组通过对应的网卡与其他服务器的网卡互联，实现各节点之间的通信。**因此相比传统网络架构，AI 服务器组网增加后端网络组网（Back End），增加了每台服务器的网络端口数量，拉动对高速交换机、网卡、光模块、光纤光缆等组件的需求。**

图23: AI 集群组网可分为前端 (Front End) 和后端网络 (Back End)



资料来源: 博通公告、开源证券研究所

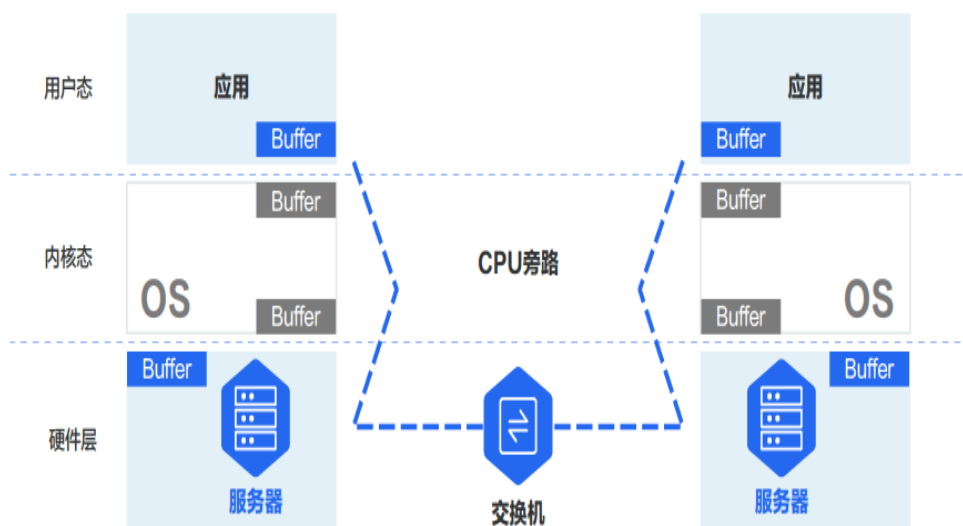
图24: 前后端网络组网均带来大量交换机需求



资料来源: Alibaba HPN

RDMA 技术广泛应用于智算中心组网。RDMA(RemoteDirect Memory Access)远程直接内存访问技术,可以绕过操作系统内核,让一台服务器可以直接访问另外一台服务器的内存,相较于传统 TCP/IP 网络,时延性能会有数十倍的改善。且有零拷贝、内核旁路、无 CPU 干预的特性,使得 AI 应用可以借助无损网络直接与远端服务器进行数据交互和内存读写,有助于消除 GPU 跨节点通信网络瓶颈,提高资源利用率和训练效率。

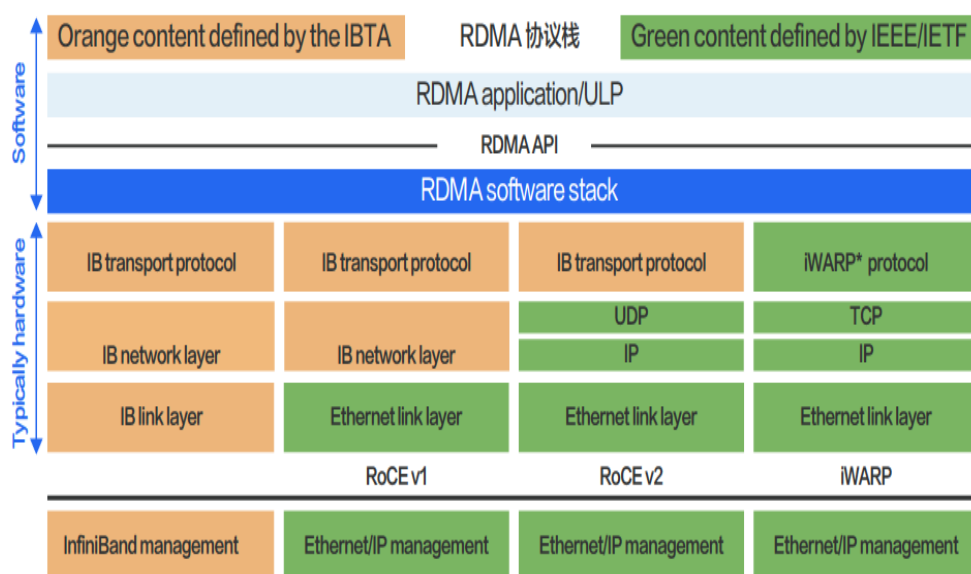
图25: RDMA 通信时延更低



资料来源: 百度智能云《智算中心网络架构白皮书》

RDMA 包含 IB、RoCE 和 iWARP 共三种实现方式。其中,iWARP 是基于 TCP/IP 的 RDMA 技术,受到 TCP 影响,性能稍差,使用较少;RoCEv1 技术当前已经被淘汰,主要以 RoCEv2 技术为主。

图26：RDMA 协议栈



资料来源：百度智能云《智算中心网络架构白皮书》

IB 网络于 1999 年由 IBTA 提出，为第一代 RDMA 技术，是一种专用无损网络，包括私有协议和专用硬件，和以太网不能互通。IB 网络从硬件层面保证数据无损，基于 credit 信令机制，确保发送端数据不会过量发送，从根本上避免缓冲区溢出分组丢失。即只有在确认下一跳有额度能接收对应数量的报文后，发送端才会启动报文发送。由于网元和网卡都必须得到授权才能发送分组，因此 IB 网络不会出现长时间拥塞，能够保证可靠传输的无损网络。

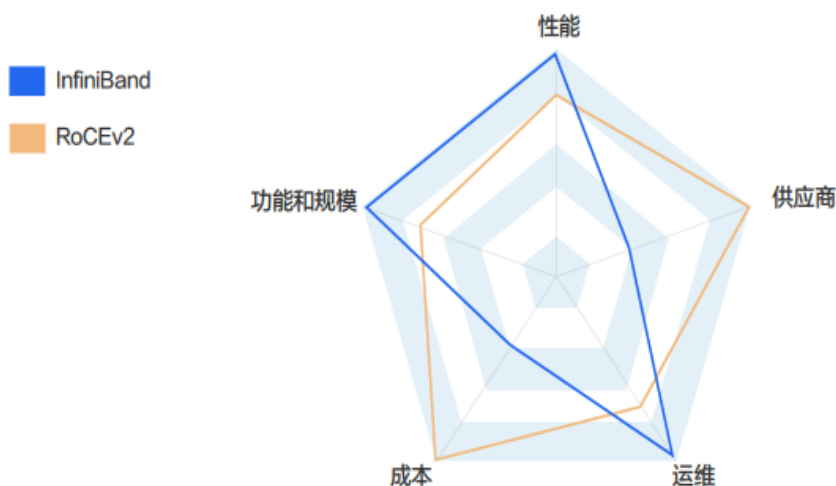
IB 网络的关键组件包括 IB 交换机、IB 专用网卡、IB 连接线缆以及子网管理器（SM），均为专用组件，与其它以太网组件不通用，IB 网络生态较为单一封闭，设备的采购和维护成本高。IB 网络设备供应商主要包括英伟达、英特尔、思科、HPE 等，其中英伟达份额最高。

RoCE 网络 2010 年同样由 IBTA 提出，基于以太网技术的传输方式，其中 RoCE v1 是链路层协议，RoCE v2 是网络层协议，支持 IP 路由，主要依靠协议如基于优先级的流量控制（PFC）、显示拥塞通知（ECN）和数据中心量化拥塞通知（DCQCN）来实现无损网络，提升可靠性。

RoCE 网络是一种纯分布式网络，将 RDMA 技术应用到传统以太网，本质上是一种网卡封装技术，只需配置支持 RoCE 的网卡和交换机即可，RoCE 网络设备的国产化厂家较多，国内数据中心交换机厂商包括华为、新华三、中兴通讯、锐捷网络等。

IB 网络时延较低，RoCE 组网性价比显著。对比 IB 和 RoCE 网络来看，（1）性能方面：IB 网络端到端时延更低，应用层性能方面更好，但 RoCEv2 性能足以满足绝大多数场景需求。（2）组网规模方面：IB 网络可支持万卡 GPU 规模集群，RoCEv2 在千卡规模集群上表现较好，组网性能仍在持续优化；（3）运维方面：IB 网络技术上更为成熟，无需复杂参数调优，部署更快。（4）成本方面：IB 组网成本较高，主要是 IB 交换机成本高于传统以太网交换机。（5）生态方面：IB 生态较为单一，以英伟达为主，RoCEv2 基于以太网，生态较为开发，供应商较多。

图27: IB 在性能和功能上更优, RoCEv2 生态丰富具备成本优势



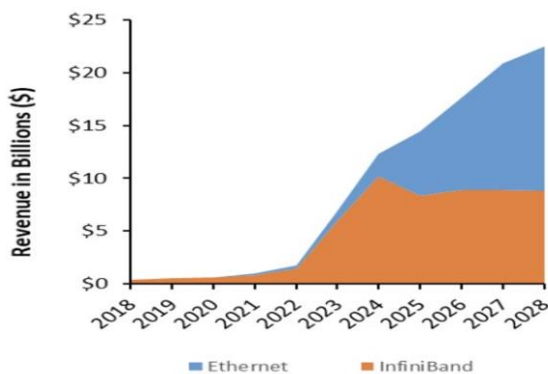
资料来源: 百度智能云《智算中心网络架构白皮书》

IB 网络仍为 AI 集群后端组网主流选择, 以太网份额持续增长。当前数据中心交换机通常用于通用服务器前端组网, 而 AI 集群中算力卡间频繁通信, AI 工作负载持续拉动服务器后端网络建设。IB 网络凭借低延迟、堵塞控制以及自适应路由等机制, 仍然主导 AI 后端网络, 但随着以太网网络部署的不断优化, 我们认为未来以太网方案占比有望持续提升, 据 650group 预测, 以太网在 RDMA 市场中占比将逐渐提升, 2027 年市场份额将超过 IB 网络, 成为主流选择:

(1) 国内外头部云厂商积极使用以太网部署 AI 网络。AWS 使用以太网为 Trainium2 GPU 组网形成 6 万以上规模的 GPU 集群; Meta 构建 2 个由 24576 个 H100 GPU 组成的算力集群, 其中一个集群使用 Arista 7800 等交换机通过 RoCE 网络组网用于 Llama3 的训练, 另一个集群采用 Quantum 2 交换机通过 IB 网络组网训练, 两种组网方式下均不存在网络瓶颈, RoCE 网络在万卡规模组网下表现良好; 字节跳动使用以太网部署万卡 GPU 的 AI 集群。

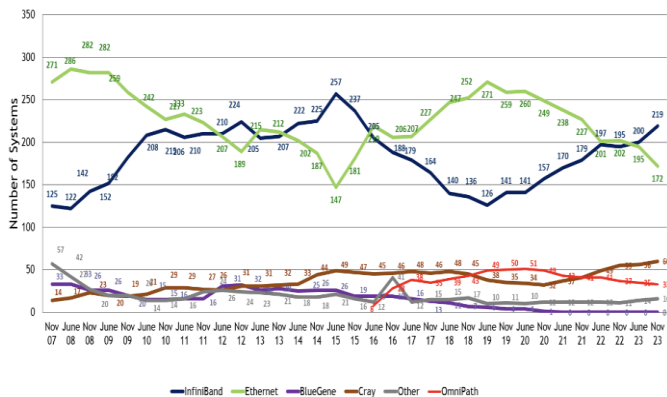
(2) 以太网根基深厚, 中腰部云厂商及大型企业或继续选用以太网。虽然当前主要需求来自头部云厂商, 但中腰部云厂商及大型企业需求仍然十分显著, 据 Dell'Oro Group 预计 Tier 2/3 CSP 厂商及大型企业需求在未来五年内接近 100 亿美元, 更偏向于使用以太网组网。

图28: 以太网 RDMA 网络市场占比有望持续提升



资料来源: 650Group

图29: HPC 集群中 IB 和以太网网络为主流方案



资料来源: Nvidia

超以太网联盟成立以对抗 IB 网络。在 AIGC 等因素催化下智算需求激增，IB 网络凭借零丢包特点在 AI 训练中独占鳌头。为取代 RoCE 协议，创建一个适用于 AI/HPC 场景、基于以太网的完整通信堆栈架构，以提高网络吞吐量、降低延迟，UEC 孕育而生，成员包括 AMD、Arista、博通、思科、华为、新华三、锐捷网络等设备商以及 Meta、微软、BAT 等云厂商。UEC 旨在优化以太网以实现高性能 AI 和 HPC 网络，最大限度地保持以太网的互操作性。我们认为，随着超以太网联盟的成立，多方厂商有望共同合作加速以太网发展，促进以太网网络在 AI 后端网络占比持续提升，逐步挑战 IB 领导地位。

图30：多方助力 UEC 发展



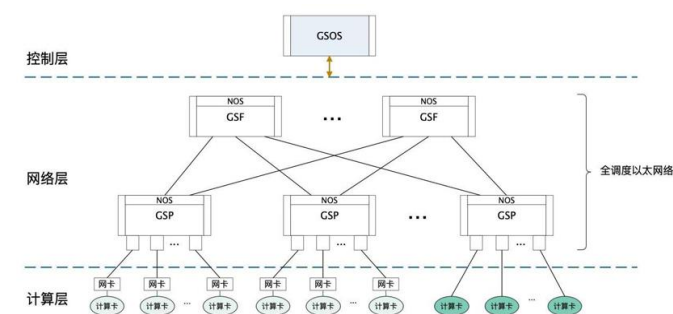
资料来源：UEC 官网

国内主导 GSE 和 ETH+协议加速智算网络建设。

2023 年 5 月，中国移动联合 10 余家企业发布全调度以太网技术架构（GSE）白皮书，并在 2023 年 8 月启动 GSE 推进计划，已有中国移动、中国联通、腾讯、华为、中兴通讯、锐捷网络、新华三、盛科通信、燧原科技等多个厂商加入。全调度以太网技术划分为 GSE1.0 和 GSE2.0 两个商用阶段，GSE1.0 基于现有芯片最大限度地支持 GSE 新技术，优化网络性能，已在中国移动智算中心（哈尔滨）超万卡集群实现首次商用，GSE2.0 则全面革新以太网底层转发机制和上层协议栈，从根本上解决传统无损以太性能和可靠性问题。

2024 年 9 月，由阿里云和中科院计算技术研究院牵头，联合平头哥、盛科通信、腾讯、字节跳动等 40 余家机构发布内首个高通量以太网（ETH+）协议标准 1.0。ETH+ 协议通过优化帧格式，实现了有效载荷比 74% 的提升；通过深度支持链路层和物理层的重传技术，ETH+以太网的语义可靠性及规模大幅提升；基于 RDMA 的在网计算技术，实现集合通信性能提升 30% 以上。

图31：GSE（全调度以太网网络）分层架构



资料来源：中国移动《全调度以太网技术架构白皮书》

图32：ETH+实现高载荷比、低时延的开放以太网

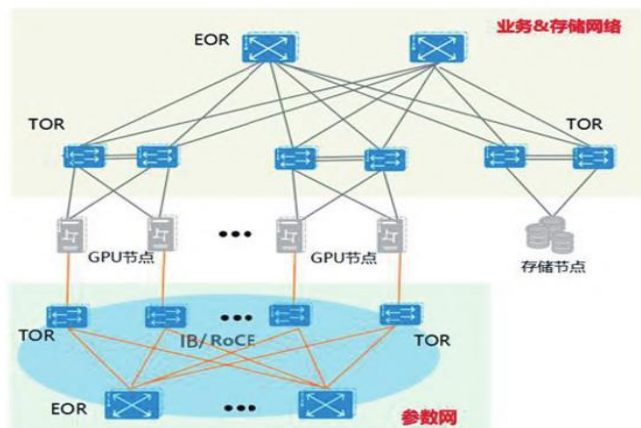


资料来源：央广网

2.2、AI/ML 后端市场快速增长，拉动交换机和网卡需求

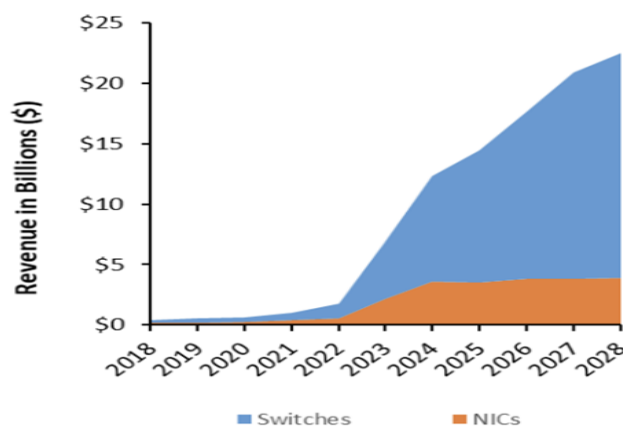
AI/ML 后端网络市场规模快速增长，拉动交换机和网卡需求。后端网络可采用运用 RDMA 技术的 RoCE 以太网和 IB 网络组网，据 650group 数据，2021 年之前，RDMA 的市场规模每年在 4 亿至 7 亿美元之间，主要受 HPC 应用的驱动。2023 年，由于 AI/ML 部署的激增，市场对 RDMA 的需求激增至 60 亿美元以上，预计到 2028 年将突破 220 亿美元，分产品来看，主要以交换机设备需求为主，分技术来看，以太网网络占比持续提升。

图33：AI 服务器网络组网架构



资料来源：李家清等《智算中心 IB 及 RoCE 网络技术探究》

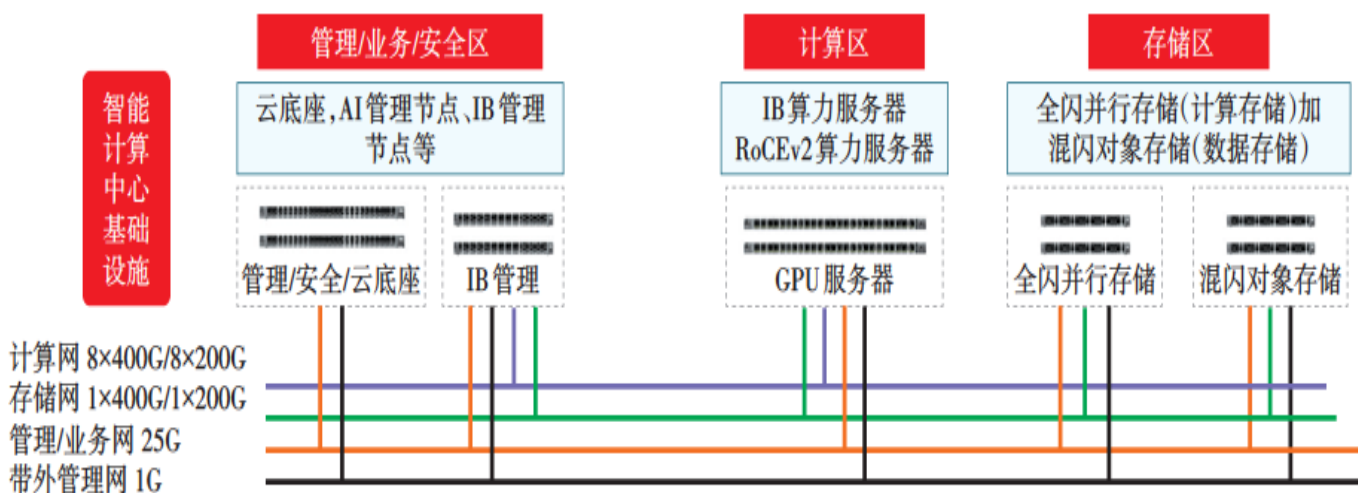
图34：RDMA 市场中交换机需求快速增长



资料来源：650 group

从各业务网络速率需求上看，计算网络需求较高。每个 GPU 对应一个高速率网络端口如 400G、800G、1.6T 等，以 SXM 8 卡 GPU 模组为例，则对应 8 个网络端口；存储网络速率需求同样较高，但端口相对较少；管理/业务网络速率则相对较低。

图35：AI 集群组网中计算和存储网络速率较高，管理和业务网络速率较低



资料来源：张世华等《新型智算中心组网方案研究》

从组网架构上看，智算 AI 集群组网需满足大带宽、无阻塞以及低时延等需求，要求数据中心交换机提供全端口线速转发的能力，并对交换机端口速率以及密度提出更高要求，交换机下联和上联带宽采用 1:1 无收敛设计，即如果下联有 32 个 800Gbps 端口，则上联也有 32 个 800Gbps 端口。

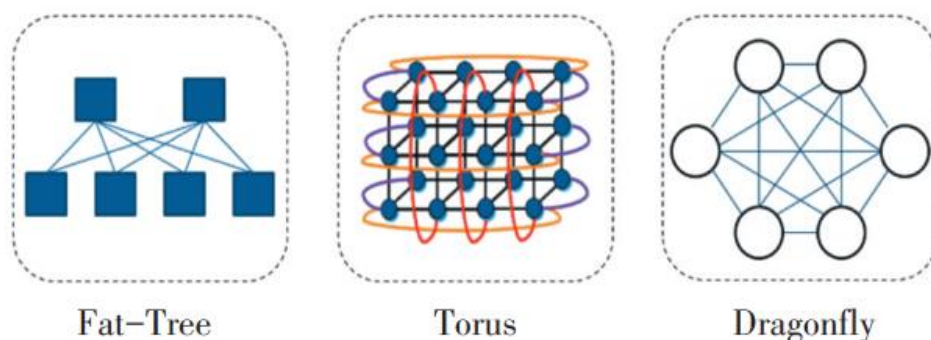
主流网络架构包含 Fat-tree、Torus、Dragonfly 三种。其中，Fat-Tree 拓扑具有网络直径短，端到端通信跳数少，建网成本低的优点，适用于中小规模智算中心。当网络达到一定规模后，例如上万节点时，可采用三层架构或改用 Dragonfly 和 Torus。Dragonfly 和 Torus 拓扑的建网成本更低，交换机端到端转发跳数明显减少，可提升网络整体吞吐和性能，适用于大规模、超大规模智算中心。

(1) Fat-Tree 是一种树形拓扑，网络带宽不收敛，支持对接入带宽的线速转发，并且在横向扩展时支持增加链路带宽。Fat-Tree 拓扑中所使用的网络设备均为端口能力相同的交换机，可有效降低网络建设成本。

(2) Torus 是一种环面拓扑，它将节点按照网格的方式排列，然后连接同行和同列的相邻节点，并连接同行和同列的最远端的 2 个节点，使得 Torus 拓扑中每行和每列都是一个环。Torus 拓扑通过从二维扩展到三维、或更高维的方式增加新的接入节点，可提高网络带宽，降低延迟。以谷歌 TPU OCS 网络为例，采用 4096 个 TPU v4 进行 3D Torus 组网。

(3) Dragonfly 是一种分层拓扑，包括 Switch、Group 和 System 3 层，其中 Switch 层包括一台交换机和与其相连的多个计算节点，交换机负责连接对应计算节点以及其他 Group 的交换机；Group 层包含多个 Switch，多个 Switch 间进行全连接；System 层包含多个 Group，多个 Group 间也进行全连接。主要优势是网络转发路径小，组网成本较低，多用在超算领域。

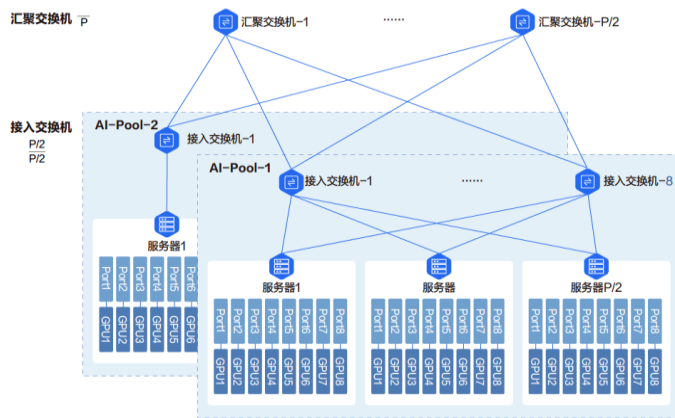
图36：主流三种智算网络架构适用不同场景



资料来源：张世华等《新型智算中心组网方案研究》

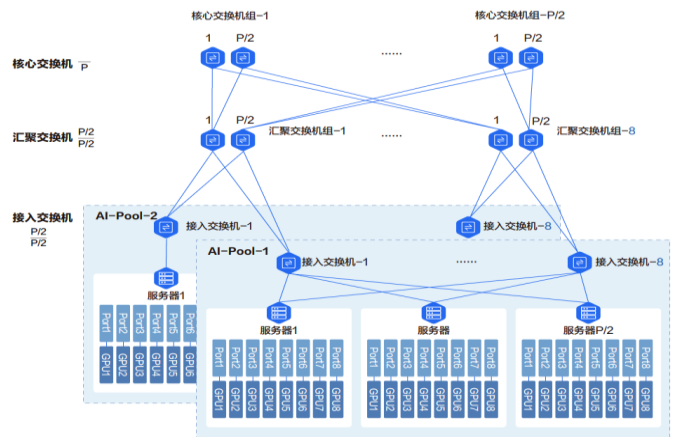
在胖树组网架构下，以搭配 8 卡 SXM GPU 模组的 AI 服务器组网为例，每个服务器 1 号网口上连至 leaf 层 1 号交换机，2 号网卡连接至 leaf 层 2 号交换机，并以此类推，直至 8 号网口连接至 8 号交换机。每 8 台 Leaf 交换机和下联的 AI 服务器组成一个 group，每 8 台 Leaf 交换机又与上面对应的 Spine 交换机组成一个 pod。若算力集群规模持续增长至 3 层组网，则以 Pod 为单位持续拓展，加入 Core 交换机进行组网，所有交换机之间均采用 Fullmesh 全连接，leaf 和 spine 层交换机上下行收敛比为 1:1 无收敛，spine 和 Core 层组网可能存在收敛比。

图37: Spine-Leaf 两层架构组网



资料来源: 百度智能云《智算中心网络架构白皮书》

图38: Fat-Tree 三层架构组网



资料来源: 百度智能云《智算中心网络架构白皮书》

两层和三层无收敛网络架构可容纳 GPU 卡规模, 取决于交换机端口数量和速率 (即交换容量=端口数×端口速率×2), 因此超大 AI 集群需要高端口密度和高速率端口的数据中心交换机。以 N 代表 GPU 卡规模, 以 P 代表单台交换机端口数量, 根据我们测算, 则两层无收敛组网架构下最多支持 $P^2/2$ GPU 卡, 对应 $3/2P$ 台交换机; 三层组网最多支持 $P^3/4$ GPU 卡。

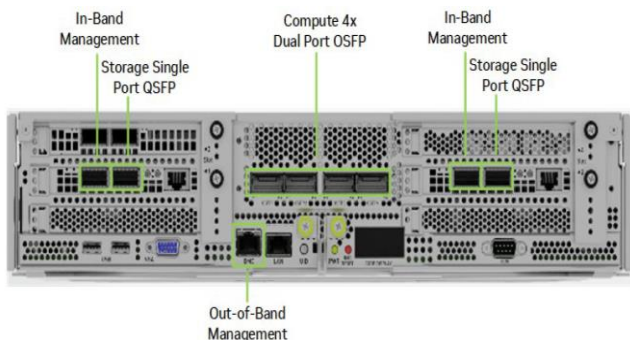
表5: 两层与三层无收敛组网的 GPU 最大节点数量

	两层无收敛组网	三层无收敛组网
最大容量	$N_{Max} = P^2/2$	$N_{MAX} = P^3/4$
GPU 卡数量		
P=40	800	16000
P=64	2048	65535
P=128	8192	524228

资料来源: 百度智能云《智算中心网络架构白皮书》、开源证券研究所

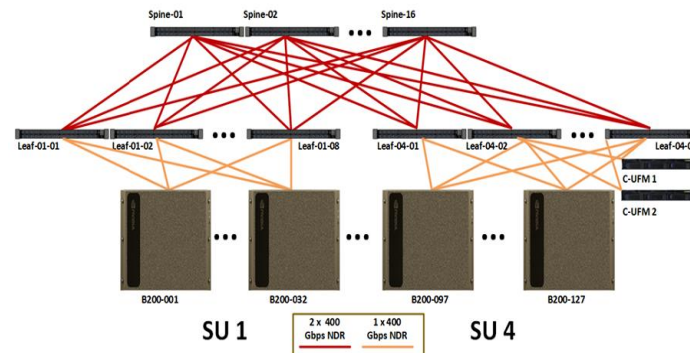
以 DGX B200 服务器、NVIDIA MQM9790 64 个 400G 端口交换机 (32 个 OSFP 端口) 为例, 服务器后端中间有 4 个双端 OSFP 对应 8 个 GPU, 31 台服务器 (即 248 个 GPU) 组成 1 个节点上联 8 台 leaf 交换机, 并对应 4 个 Spine 交换机, 共计 12 台计算网络节点交换机。此外, 存储网络仍需配套高速 400G 交换机, 管理网络则速率较低 (100Gbps), 均会带来大量数据中心交换机需求。

图39: DGX B200 网络端口示意图



资料来源: Nvidia

图40: DGX B200 127 节点计算网络组网架构



资料来源: Nvidia

表6: B200 组网服务器与交换机对应关系

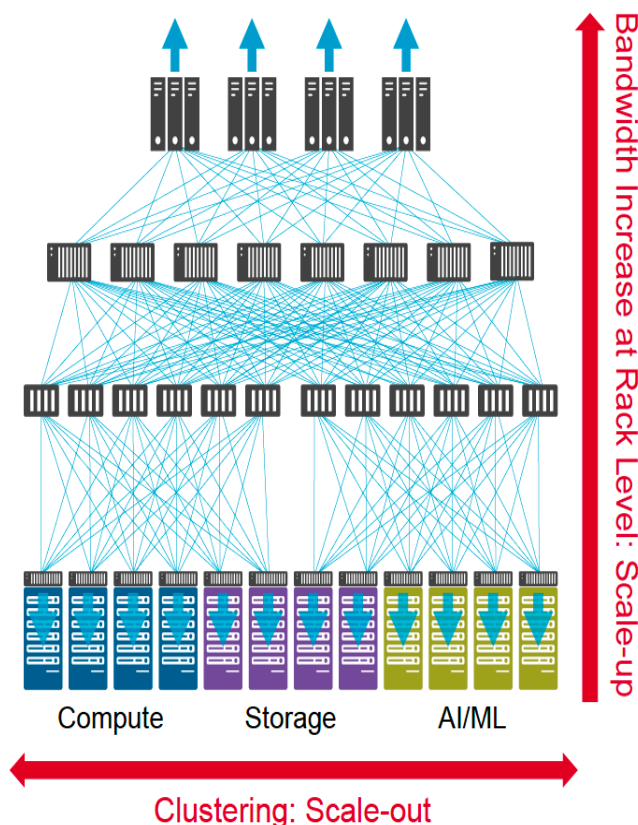
计算节点 (个)	服务器数量 (个)	GPU 数量 (个)	IB 交换机数量 (个)		网线数量 (条)	
			Leaf	Spine	Leaf-服务器	Spine-Leaf
1	31	248	8	4	252	256
2	63	504	16	8	508	512
3	95	760	24	16	764	768
4	127	1016	32	16	1020	1024

数据来源: Nvidia、开源证券研究所

2.3、未来组网架构—Scale up 和 Scale out 的探讨

为支持万亿或更大参数量模型持续发展，未来集群规模或将从万卡逐渐向十万卡或更大规模扩展，对于超节点及超大规模组网架构，未来有望从 Scale up 和 Scale out 两个维度来实现总算力规模的提升。

图41：算力集群拓展方向 Scale up + Scale out

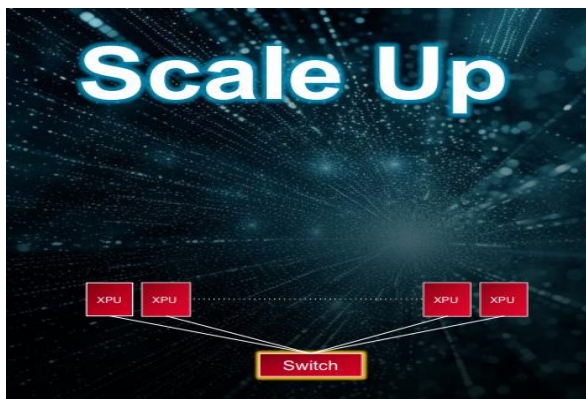


资料来源：博通公告

Scale up：主要通过提高单个节点内的算力规模，进而提升集群的算力规模。在服务器层面增加算力芯片总数，以 A100、H100、B200 DGX 系列为例的单个 AI 服务器内部算力模组主要由 8 张算力卡内部通过 NVSwitch 芯片互联组成，未来有望通过引入支持更多算力芯片互联比如 16 卡、32 卡互联的 Switch 芯片，以优化 GPU 南北向的互联效率和规模，增强张量并行或 MoE 并行的数据传输能力，同时提升 GPU 卡间互联带宽，通过高速互联总线将更多算力芯片互联，提升单服务器算力性能；在机柜层面增加服务器总数，以 GH200 NVL32、GB200 NVL72 为例，单机柜内部通过引入更多服务器再搭配高速交换机实现互联，提升单机柜算力性能，再通过机间互联扩展至 NVL576，提升单个节点的算力性能。

Scale out：主要通过高速互联容纳更多节点，进而提升集群整体算力规模。当前机间通信主要以 400G/800G 为主，未来有望通过更高速率如 1.6T 组网互联，以提高互联带宽，支持更多节点高速互联；采用 CPO (Co-Packaged Optics) /NPO (Near Packaged Optics)、多异构芯片 C2C (Chip-to-Chip)封装等方式降低延时，进而提升数据传输效率；通过增加交换机端口数量提升相同架构下的 GPU 节点数量上限，或通过增加集群组网规模以实现更多节点间互联，如从 2 层胖树组网增加至 3、4 层组网架构，或改由 Torus、Dragonfly 等方式组网，实现从千卡向万卡、十万卡集群拓展。

图42: Scale up——更多 GPU 互联



资料来源：博通公告

图43: Scale out——更多节点互联



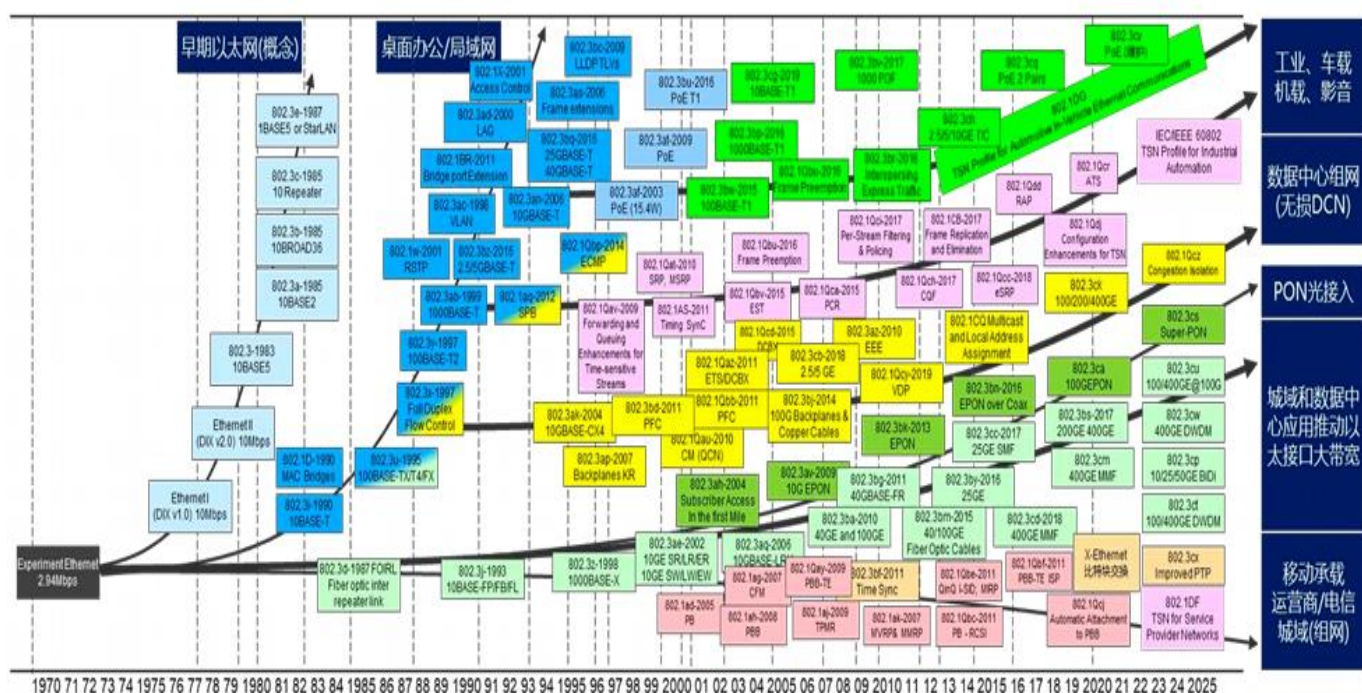
资料来源：博通公告

3、交换机变革 2：800G 交换机开始放量，102.4T 交换芯片有望推出

3.1、AI 大幅提升算力需求，驱动以太网交换机需求增长

以太网的起源可以追溯到 1973 年，梅特卡夫发明了基于 Aloha 网络的新系统，改进了 Aloha 可随意访问共享通信信道的机制，能够把任何计算机连接起来，实现计算机之间的数据传输，该系统被其命名为以太网。3 年后，以太网局域网时代正式开始，为了能接入更多不同的设备，以太网技术走上标准化之路。此后，以太网进入了高速发展的 40 年，网络速度快速提升，从 10Mbps 到百兆、再到千兆、万兆以太网，到现在以太网已具备 400、800Gbps 的商用能力；应用范围也从最初的局域网，进入到城域网和广域网。

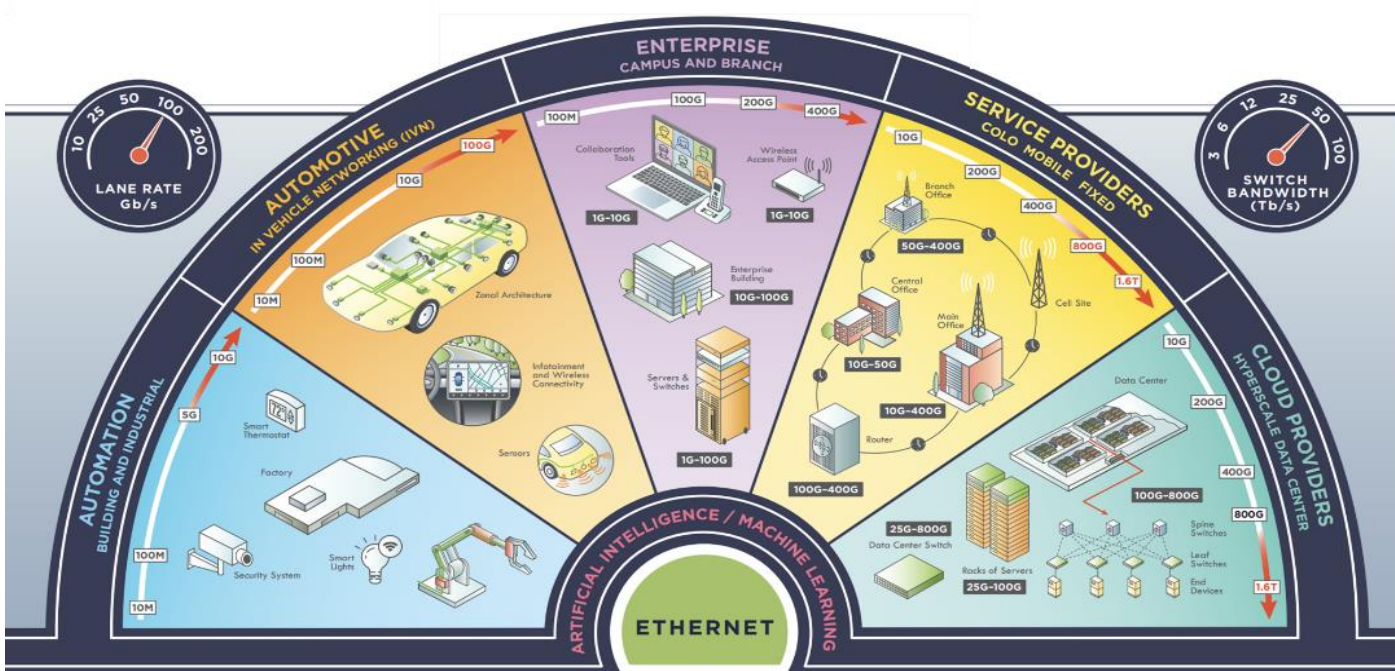
图44：数据中心以太网接口带宽向 800G、1.6T 发展



资料来源：ODCC《下一代以太网技术需求白皮书》

经过几十年的发展，以太网已经广泛运用在生活中的多个场景如汽车、工业、企业和校园、运营商和云服务网络等。分场景来看，（1）汽车向智能化、网联化发展，车内高速以太网在丰富的多媒体需求以及自动驾驶辅助系统（ADAS）等需求推动下高速发展，以太网速率从 10M 逐渐向 100G 提升；（2）传统工业持续推进数字化转型，工业 4.0、新型工业化背景下，工业生产数据互联互通，从传统的现场总线网络转向以太网，以太网速率从 10M 逐渐向 10G 或更高速率提升；（3）运营商网络多年来持续推动以太网迭代，在 DCI、PON 光接入、OTN 等多个场景中运用以太网技术，随着 5G-A、6G 时代来临，以太网有望向 800G、1.6T 更高速率升级；（4）AIGC 浪潮下，云服务厂商加速部署高速高密度网络，随着 AI 模型参数持续增长，算力节点之间的互联带宽需求高速增长，持续推动以太网速率扩展 800G、1.6T。

图45：工业自动化和汽车场景以太网速率较低，运营商和云计算网络以太网速率需求较高



资料来源：Ethernet Alliance

由于交换机使用场景较为分散，行业产业链参与公司众多。交换机产业链上游主要包括芯片、元器件、光模块、电路板、网络操作系统、电源模块和结构件等元件；中游按照终端应用场景，可分为工业交换机、运营商交换机、数据中心交换机、园区交换机等；下游应用于电信运营、云服务、数据中心等领域。从交换机出货形式可分为传统交换机、白盒交换机和裸金属交换机，其中，传统品牌交换机厂商主要包括思科、华为、新华三、Juniper、中兴通讯、Mellanox 等，白盒交换机厂商主要包括 Arista、锐捷网络、新华三等，裸金属交换机包括 Accton、Quanta、Alpha Networks 等。交换机作为数据中心的网络底座，随着数据中心的持续建设，有望带动数据中心交换机的需求。

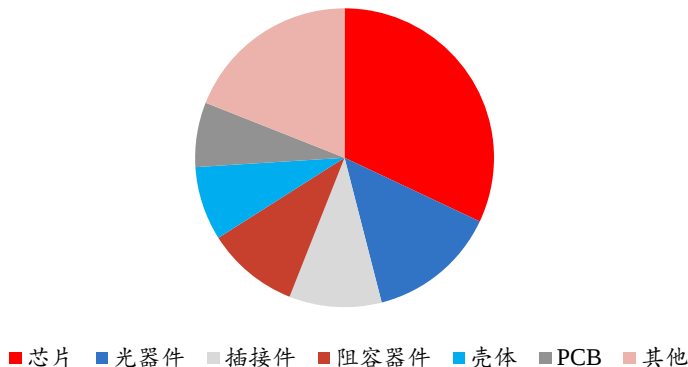
图46：交换机产业链涉及厂商众多



资料来源：中商情报网、各公司官网、开源证券研究所

从交换机原材料成本结构来看，芯片成本占比达到 32%，包含以太网交换芯片、CPU、PHY、CPLD/FPGA 等，其中以太网交换芯片和 CPU 是最核心部件，其次为光器件、接插件、壳体、PCB 等。

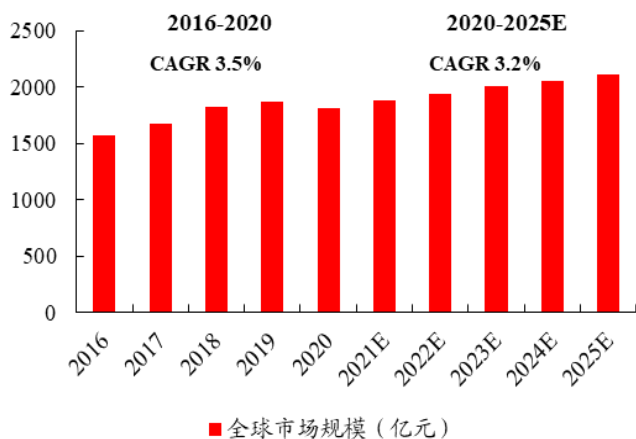
图47：芯片类在以太网交换设备的成本占比最大



数据来源：亿渡数据、开源证券研究所

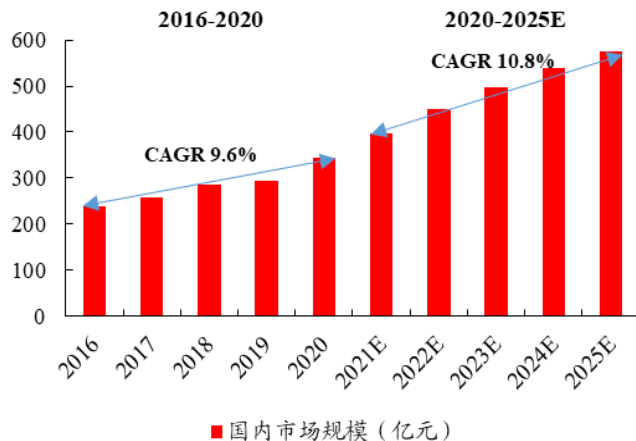
国内以太网交换设备市场处于快速发展阶段。全球以太网交换设备市场发展较为成熟，根据 IDC、灼识咨询数据预测，2025 年全球以太网交换设备市场规模有望达到 2112 亿元，2020-2025 年的年复合增长率预计将达到 3.2%。国内以太网交换设备市场处于快速发展阶段，灼识咨询数据显示，2020 年国内以太网交换设备市场规模为 343.8 亿元，预计 2025 年市场规模有望达到 574.2 亿元，2020-2025 年均复合增长率为 10.8%，年均增速大约是全球市场的三倍，国内市场在全球市场的比重有望从 2020 的 19% 提升至 2025 年的 27.2%，占比实现大幅提升。

图48：预计全球以太网交换设备市场 2020-2025 年复合增速为 3.2%



数据来源：IDC、灼识咨询、开源证券研究所

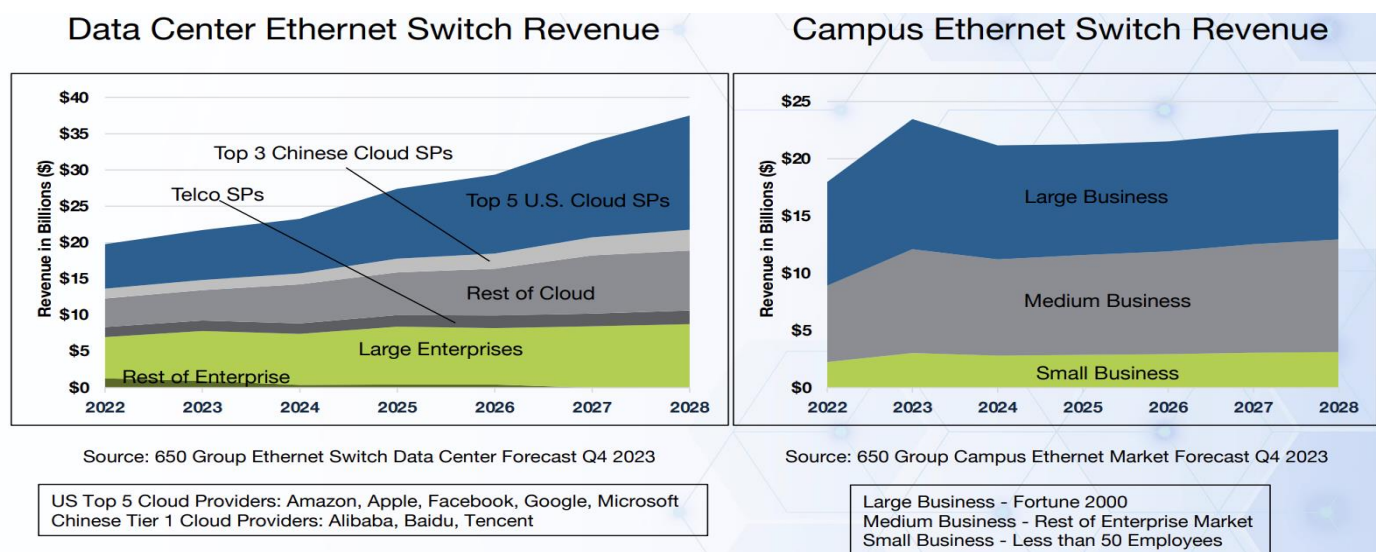
图49：预计国内以太网交换设备市场 2020-2025 年复合增速为 10.8%



数据来源：灼识咨询、开源证券研究所

数据中心以太网交换机主要客户为云厂商，园区交换机主要客户为中大型企业。以太网交换机市场可分为数据中心交换机市场、园区交换机市场、运营商市场和工业交换机市场，其中，数据中心和园区交换机市场空间较大。据 650 Group 数据，预计 2024 年全球数据中心和园区以太网交换机市场规模都将超过 200 亿美元，从下游客户结构来看，数据中心交换机客户主要以北美五大云厂商为主，园区交换机客户主要为中大型企业。

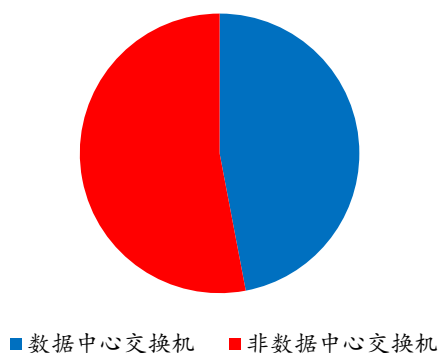
图50：数据中心以太网交换机下游客户主要是云厂商和大型企业，园区交换机客户主要为中大型企业



资料来源：650 Group、Arista network 2024Q1 财报

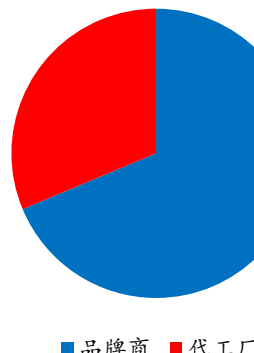
国内数据中心交换机占比有望持续提升。据中商情报网数据，国内交换机市场以数据中心交换机为主，2022 年，国内数据中心用以太网交换机收入占比达 47%，我们认为园区交换机需求受宏观经济发展节奏影响较大，当前需求增长放缓，相反，AIGC 发展迅猛，或将带动数据中心交换机持续放量，占比有望持续提升。从交换机制造商结构上来看，仍以品牌商自产为主，2021 年国内交换机品牌商占比 68.7%。

图51：2022 年国内数据中心交换机占比接近一半



数据来源：中商情报网、开源证券研究所

图52：2021 年国内交换机制造以品牌商为主

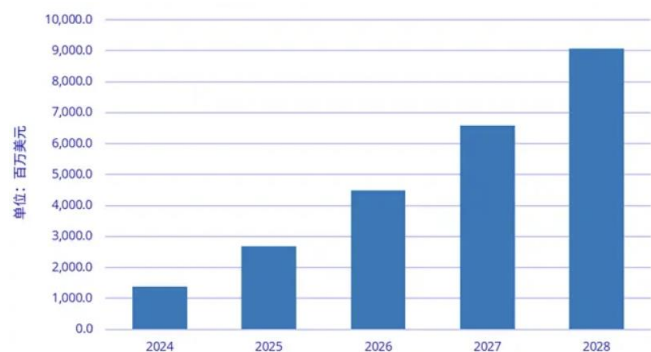


数据来源：中商情报网、开源证券研究所

随着 AIGC 持续发展，交换机作为算力网络底座，需求有望加速释放。据 IDC 预测，全球生成式 AI 数据中心以太网交换机市场将以 70% 的年复合增长率快速增长，有望从 2023 年的 6.4 亿美元增长到 2028 年的 90.7 亿美元。

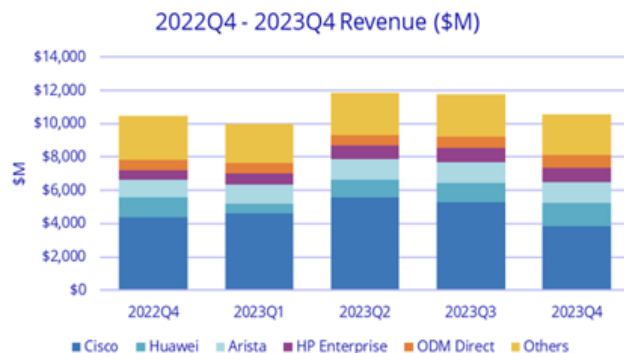
全球以太网市场份额较为集中，CR5 超 75%，思科份额领先，华为、新华三份额靠前。交换机从全球市场份额来看，2023 年，思科以太网交换机营收同比增长 22.2%，非数据中心交换机占比 69.5%，仍居以太网交换机市场首位，市场份额达到 43.7%；Arista 凭借数据中心交换机放量，以太网交换机营收同比增长 35.2%，市占率达到 11.1%；华为以太网交换机营收同比增长 10.6%，市场份额达到 9.4%；HPE 以太网交换机营收同比增长 67.6%，其中 89.6% 来自非数据中心领域，市占率达到 9.4%；新华三 2023 年市占率达到 4.2%。

图53：全球生成式 AI 数据中心以太网交换机市场规模有望快速增长



资料来源：IDC

图54：2023 年全球以太网市场份额较为集中

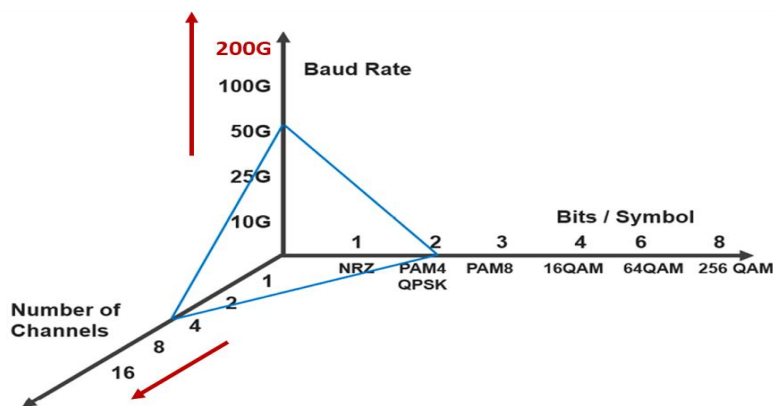


资料来源：IDC

3.2、交换芯片不断升级，102.4T 芯片有望于 2025 年底推出

通常有三种方式来增加数据中心以太网互联速率：（1）使用更复杂的信号调制技术，以提高比特速率，增加传输效率，如 PAM4 信号调制技术的比特速率是 NRZ 信号技术的 2 倍；（2）增加单通道速率或波特率；（3）增加通道数量。

图55：三种方式提升互联速率

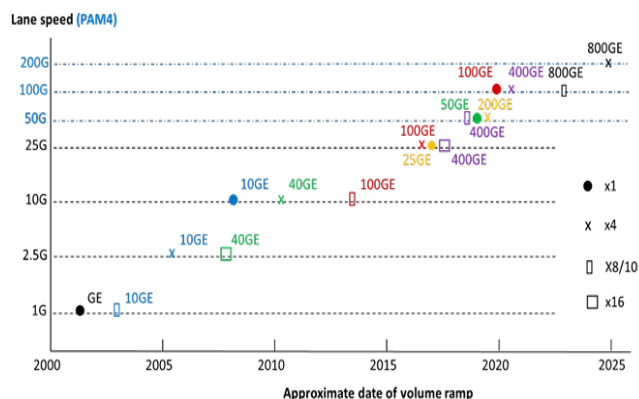


资料来源：Keysight、开源证券研究所

交换芯片平均迭代周期约为 2 年，下一代 102.4T 交换芯片有望于 2025 年底推出。数据中心网络对高性能交换机的需求推动以太网交换芯片的飞速发展，以博通开发的数据中心交换芯片 Tomahawk 系列为例，第一代 Tomahawk 芯片于 2014 年下半年发布，带宽 3.2Tbps，采用 25Gbps SerDes 技术，支持 32 个 100G 端口；2022 年下半年，Tomahawk5 发布，单芯片带宽高达 51.2Tbps，采用 112Gbps SerDes 技术，支持 64 个 800G 端口（单芯片最多具有 64 个集成 Peregrine SerDes 内核，每个内核集成 8 个 SerDes 和相关 PCS）。

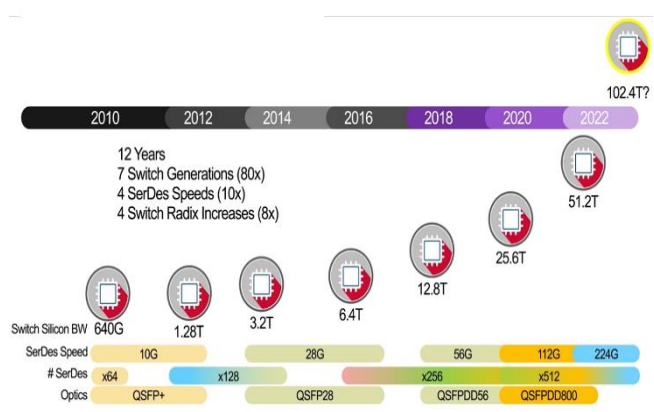
芯片制程由 2014 年的 3.2T 22nm 快速演进至 2022 年的 51.2T 5nm，据博通公开电话交流会，下一代 102.4T 芯片有望于 2025 年底推出，或将采用 3nm 制程，单芯片功耗存在超过 1000W 的可能，或切换至液冷散热模组，我们认为下一代芯片或将沿用 PAM4 技术，SerDes 速率或将达到 224Gbps，通道数量保持 512 个，并且延时方面更低以支持 AI 集群网络发展。

图56：通道速率向 224Gbps 提升



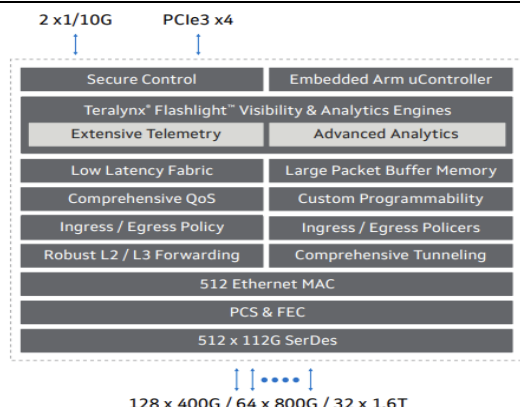
资料来源：Earlwood Marketing

图57：交换芯片迭代周期约为 2 年



资料来源：Keysight

图58：美满 Teralynx10 51.2T 芯片支持 32 个 1.6T 端口



资料来源：Marvell 公告

图59：新华三首发 1.6T 端口智算交换机



资料来源：新华三公众号

AI 高密度训练需求下，高交换容量交换机需求持续增长，采用多芯片盒式交换机的形式有望填补芯片迭代真空期带来的盒式交换容量瓶颈。2024 年 3 月，英伟达在 GTC 大会上发布 Quantum-X800 系列交换机，包含 4 颗交换芯片，可实现端到端 800Gb/s 吞吐量。以 Q3400-RA 型号为例，整体高度 4U，可实现 144 个 800G 端口分布在 72 个 OSFP 端口中，总交换容量带宽达到 115.2Tbps，单个隧道为 200Gb/s SerDes，其中，Q3400 仍采用风冷散热设计，Q3400-LD 采用液冷散热设计。由于单交换机包含 4 颗交换芯片，交换机容量增长带动可支持高速率端口数量增长，可充分满足 AI 集群的高密度组网需求，两级胖树拓扑结构下，可连接至多 10368 个 NIC 网卡，Quantum-X 以太网系列已被 Azure 和 Oracle 云采用。

图60：英伟达多芯片盒式交换机包含 72 个 1.6T OSFP 端口



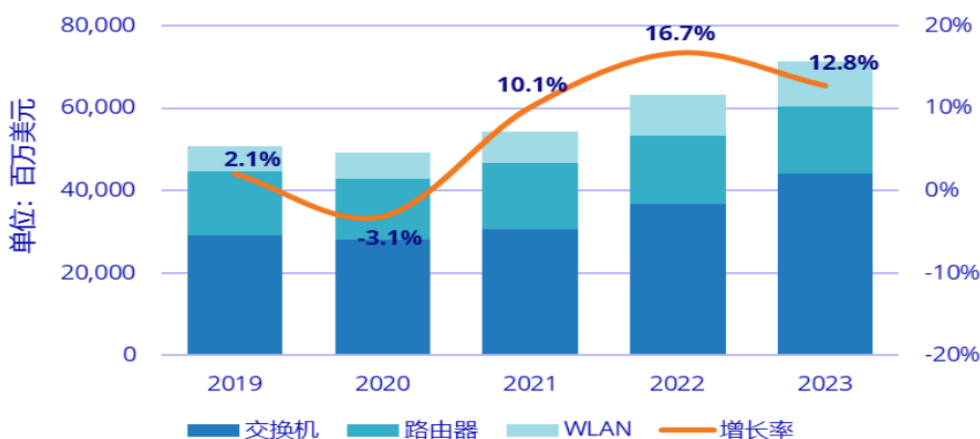
资料来源：Nvidia 公告

中国移动主导 GSE 芯片研发，国产 51.2T 交换芯片加速发展。2024 年 9 月，中国移动启动 GSE 芯片合作伙伴招募，计划向 GSE 交换芯片方向投入上亿元资金，与合作伙伴共同开发一款高规格的（51.2T 以上），适用于智算、通算、超算等场景的芯片产品，国产交换芯片有望加速追赶。

3.3、AI 拉动高速交换机需求，全球 800G 交换机开始放量

AIGC 持续带动数据中心市场持续增长，800G 端口数据中心交换机有望于 2024 年开始放量，400G 需求加速释放。据 IDC 数据，2023 年全球以太网交换机市场规模达到 442 亿美元，同比增长 20.1%，全球企业及运营商路由器市场规模达到 164 亿美元，同比基本持平。分市场结构来看，2023 年数据中心市场规模达到 183 亿美元，同比增长 13.6%，占比达到 41.5%，其中，100G 端口交换机仍为市场主流，占数据中心市场 46.3%，营收同比增长 6.4%，200/400G 端口交换机营收同比增长 68.9%，ODM 厂商直销营收达到 63 亿美元，同比增长 16.2%，占数据中心市场 14.3%；非数据中心交换机市场规模达到 259 亿美元，同比增长 25.2%，其中，1G 端口交换机仍占主流，占非数据中心市场 56.5%，同比增长 24.2%，10G 端口交换机 2023Q4 占比 20.4%，全年营收同比增长 5.3%。分地区来看，美国地区市场 2023 年同比增长 28.8%，中国市场虽然全年下跌 4.0%，但在 2023Q4 同比增长 9.1%，市场需求持续回暖。

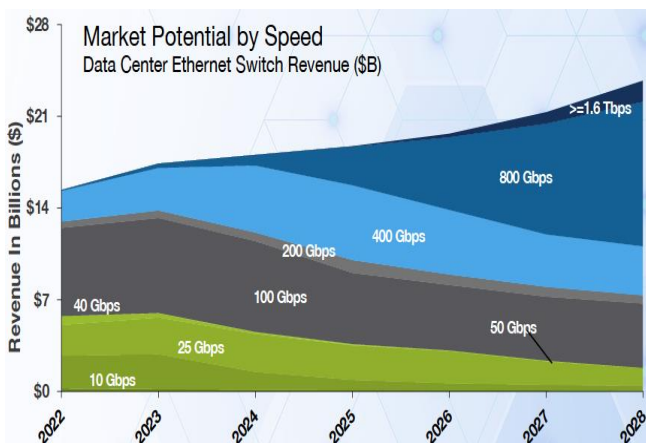
图61：2023 年全球以太网交换机市场持续增长



资料来源：IDC

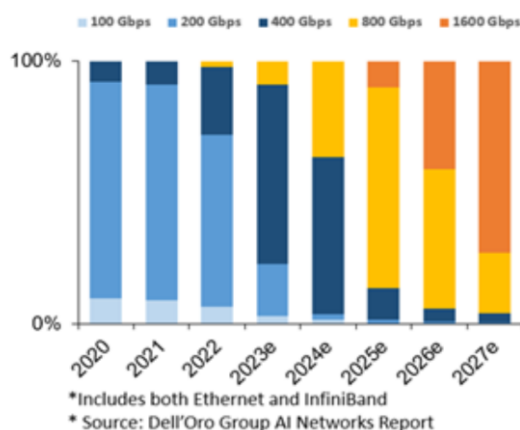
据 Dell'Oro 数据，从端口速率来看，2023 年全球 100G 端口数据中心交换机仍为主流，400G 端口交换机加速放量，预计 2024 年 800G 端口交换机有望逐渐放量，并逐渐成为主流，1.6T 端口交换机有望于 2026 年左右开始放量。对于 AI 后端网络，Dell'Oro 预计到 2027 年，AI 后端网络中几乎所有端口都将以 800 Gbps 的最低速度运行，其中 1600 Gbps 占端口的一半，网络带宽将以 3 位数复合增速迅速提升。

图62：2024 年 800G 端口交换机有望加速放量



资料来源：Dell'Oro，注：预测不包含以太网 AI

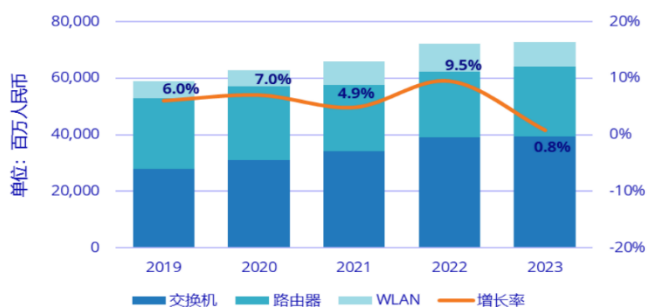
图63：AI 后端网络速率有望加速迭代



资料来源：Dell'Oro

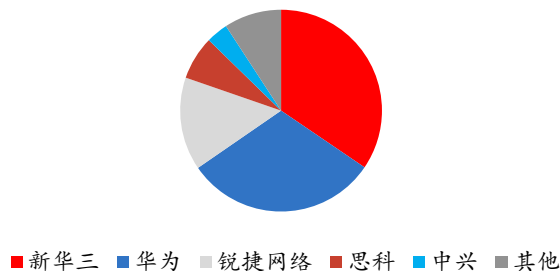
国内数据中心交换机市场持续增长，加速向 400G 端口交换机迭代。据 IDC 数据,2023 年中国交换机市场规模同比增长 0.7%,其中数据中心交换机同比增长 2.2%。随着 AI 模型的快速发展,数据中心超大规模组网需求持续提升,网络需求由云数据中心 CPU 计算的 10G~100G 上升至 GPU 训练的 100G~400G,预计 2024 年 400G 端口出货量将继续增长,51.2Tb 芯片的成熟商用也将助推 400G 端口的采用。2023 年园区交换机市场受到宏观经济波动影响,同比下滑 0.5%,伴随宏观经济好转,园区交换机部署节奏有望回到正轨。

图64：2023 年国内数据中心交换机市场持续增长



资料来源：IDC

图65：2023Q1 国内新华三、华为、锐捷网络份额靠前

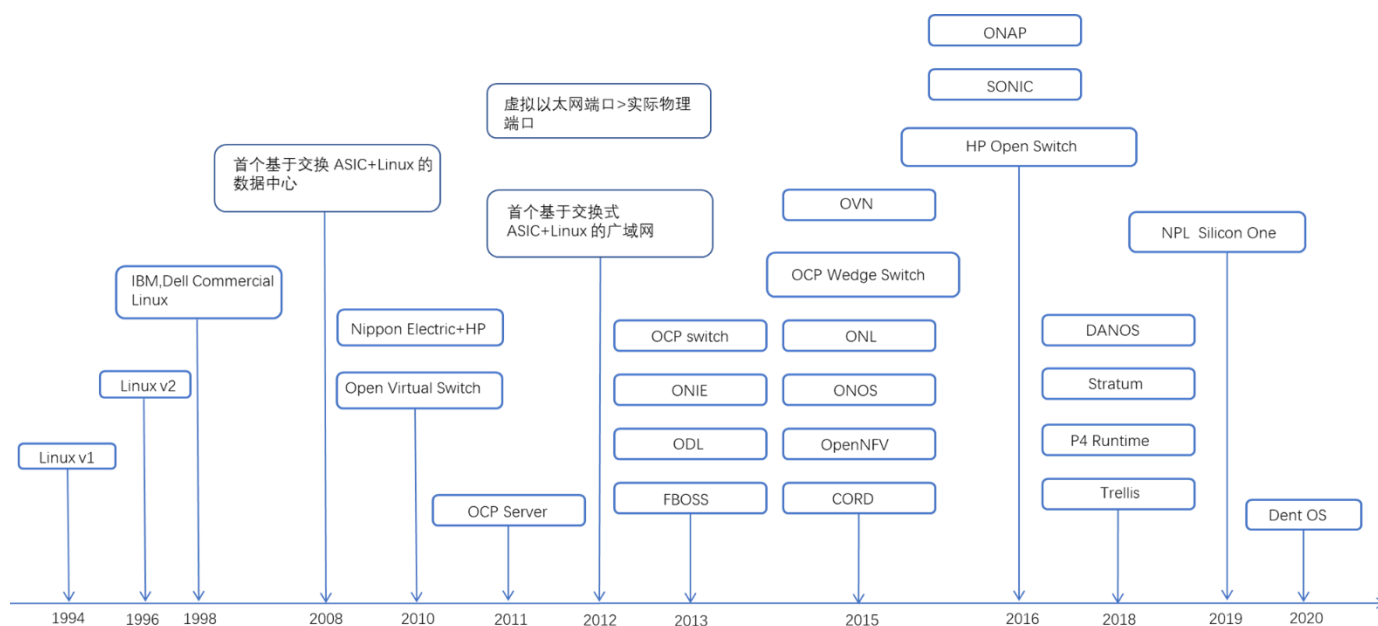


数据来源：IDC、开源证券研究所

4、交换机变革 3：交换机白盒化趋势显著，带来新成长机遇

多方助力，白盒交换机在过去 30 年间加速发展。1994 年，Linux 1.0 版本正式发布，2 年后 2.0 版本正式更新，提供了网络协议/功能控制的开源框架。用户可根据自己的需求，对网络功能与协议进行修改和定制。2013 年，OCP 开启交换机硬件白盒化的标准化工作，2015 年，第一款白盒交换机 Wedge 正式亮相。至今，白盒设备、软件操作系统、网络自动化等技术蓬勃的发展，白盒交换机生态不断完善。

图66：白盒交换机在过去 30 年间加速发展

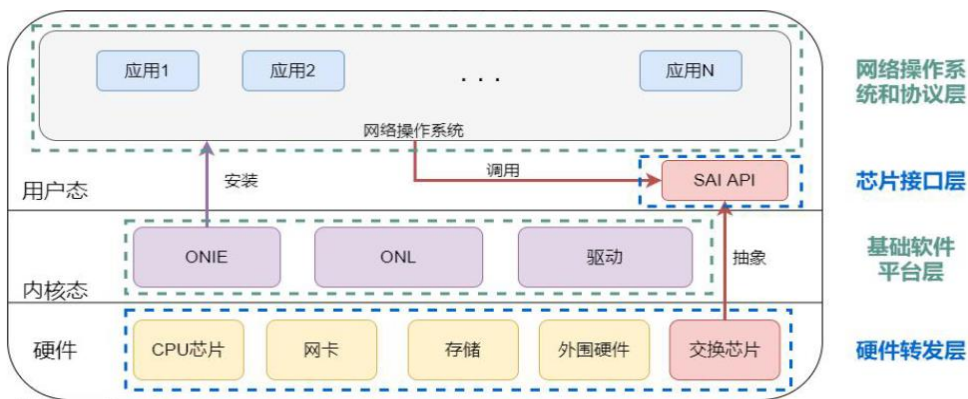


资料来源：网络通信与安全紫金山实验室《未来网络白皮书》、张成林等《面向未来网络的白盒交换机体系综述》、开源证券研究所

白盒交换机核心在于解耦。传统交换机采用软硬一体化设计，底层芯片与上层系统紧密捆绑，白盒交换机是一种硬件与软件解耦的网络交换机，利用标准化芯片接口解耦底层芯片和上层应用，其硬件由开放化的硬件组件组成，而软件（包括操作系统和网络功能等）可以由用户或第三方自由选择和定制。而裸金属交换机只包含硬件，由用户自主购买或者选择软件操作系统。

从硬件上来看，主要包括（1）交换芯片：用于交换转发数据包，是交换机的核心部件，白盒交换机芯片要求接口标准化，解耦底层芯片和上层应用；（2）CPU 芯片：主要管控系统运作；（3）网卡：提供 CPU 侧管理功能；（4）存储器件：包括内存、硬盘等；（5）外围硬件：包括风扇、电源等，接口、结构等需要符合 OCP 或其他标准化规范。从软件上看，软件主要指网络操作系统（NOS）以及其所搭载的网络应用，NOS 一般通过基础软件平台的引导来完成安装，而芯片接口层会将交换芯片的硬件功能封装成统一的接口，从而实现上层应用与底层硬件的解耦。

图67：白盒交换机硬件与软件解耦



资料来源：网络通信与安全紫金山实验室《未来网络白皮书》

白盒交换机灵活性、可扩展性较高。白盒交换机不同于传统品牌交换机，相比传统交换机，白盒交换机灵活性、可扩展性较高，采购和维护成本较低，广泛应用于互联网和运营商网络。白盒交换机产业生态较为完善，上游主要为硬件提供商包括 Arista、思科、新华三、锐捷网络、Accton、工业富联、Dell、Quanta 等，网络操作系统供应商包括 Arrcus, Kaloom, Cumulus, Big Switch、FBOSS、SONIC 等，下游客户主要包括云服务商、电信运营商等，主要利用白盒交换机用于业务转型和网络重构。

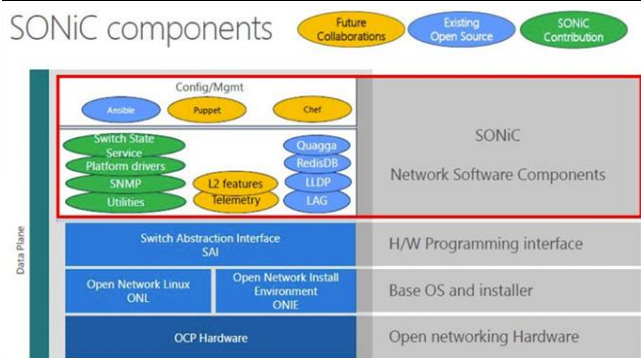
图68：白盒交换机产业生态较为完善



资料来源：各公司官网、网络通信与安全紫金山实验室《未来网络白皮书》、SDNLAB、开源证券研究所

Sonic 逐渐成为超大数据中心网络首选开源系统，白盒交换机市场空间持续增长。2016 年，微软在 OCP 峰会上正式发布 SONiC(Software for OpenNetworking in the Cloud)开源交换机操作系统，SONiC 将网络软件与底层硬件分离，并建立在交换机抽象接口（SAI）的 API 之上，SAI 为 ASIC 提供统一接口。目前，SONiC 作为一个成熟的构建交换机网络功能的软件集架构，可实现数据控制面与转发面的分离，用户通过购买白盒交换机搭载 SONiC 实现不同网络功能，能够更快调试、测试、更改软件策略和拓扑，进而实现新的网络架构，已被 BBAT、微软、谷歌等国内外多个云厂商规模部署运行，其中大部分为单芯片盒式交换机。据 Omdia 数据，2022 年数据中心以太网交换机端口出货量增长 12%，其中思科市场份额为 37%，Arista 18%，华为 8%，H3C 7%，白盒供应商份额占比 14%，同比增长 4 个百分点。

图69: SONiC 软件系统结构图



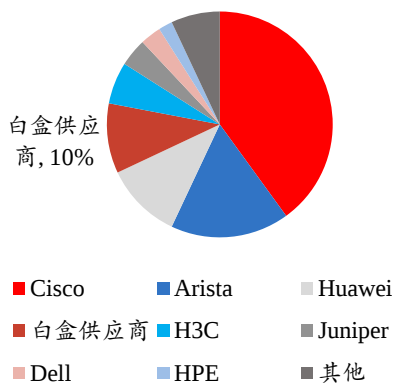
资料来源: ODCG《框式开放自研交换机技术实现与应用场景白皮书》

图70: SONiC 参与厂商众多



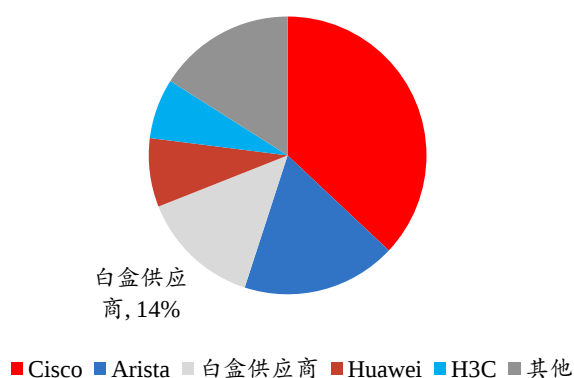
资料来源: Sonic 官网

图71: 2021 年白盒交换机端口出货份额占比 10%



数据来源: Omdia、开源证券研究所

图72: 2022 年白盒交换机端口出货份额占比 14%



数据来源: Omdia、开源证券研究所

5、交换机变革 4：光交换机逐渐成熟，光电融合组网落地大模型训练

光交换机可靠性更强，功耗更低。在光电融合交换方案中,光交换功能模块的主要方案分为光电路交换（OCS）、光突发交换(OBS)和光包交换(OPS)3种。

光电路交换机（OCS）主要通过配置光交换矩阵,从而在任意输入/输出端口间建立光学路径以实现信号的交换。

（1）由于光传播路径的宽带和无源的特性,OCS 对光信号的速率和协议等均是透明的,不需要随着服务器 NIC 网卡速率以及端口迭代，相同 OCS 硬件可以跨代际的被重复利用，长期成本开支更低，生命周期较长；

（2）由于没有光/电转换和相应的包处理和分发的过程,OCS 拥有更小的每端口功耗，以 400Gbps 端口为例，OCS 每端口功耗<1W ;电交换机每端口功耗>10W)，以及较低的时延（OCS 时延数十 ns，EPS 时延百 μ s)；

（3）由于 OCS 整机使用芯片类型及数量较少，故障率远低于电交换机，可靠性更强。

由于 OCS 缺乏包处理能力，只能将某个输入/输出端口连通配对，只有当全局动态的流量预测与实时光交换矩阵配置完美结合时，OCS 才能较好满足业务需求，而传统业务流量通常难以预测，成为了制约 OCS 规模应用的重要因素，但 AI 大模型预训练基于已知的数据集和模型算法，具有流量可预测的特点，进而催生了 OCS 的众多应用形式，当前光电融合方案中 OCS 方案商用化程度较高，基于 3D-MEMS 系统的 OCS 方案综合应用较好。

表7：WDM/MEMS 方案商用程度较好

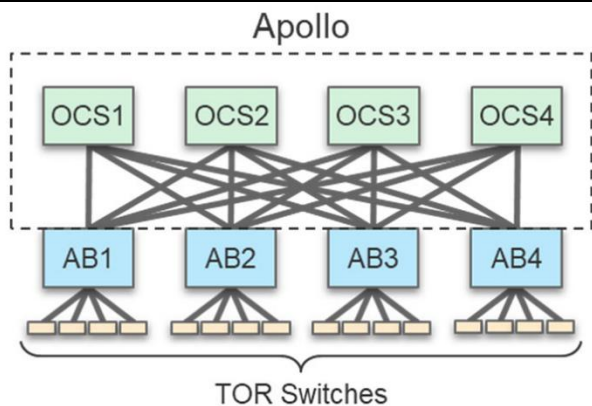
架构名称	交换类型	控制方式	光学设备	连接方式	可扩展性	商用化水平
Helios	EPS/ OCS	分布式	WDM	光纤	中	商用设备
C-Through	EPS/ OCS	集中式	MEMS	光纤	低	商用设备
Mordia	OCS	集中式	WSS	光纤	低	实验原型
RotorNet	OCS	分布式	MEMS	光纤	高	实验原型
Sirius	OCS	分布式	AWGR、TunableLaser	光纤	高	实验原型
Proteus	OCS	集中式	WSS、MEMS	光纤	低	商用设备
OSA	OCS/ OPS	集中式	WSS、MEMS	光纤	低	商用设备
Lotus	OCS/ OPS	分布式	AWGR	光纤	中	软件仿真
ReSAW	OPS/ OBS	分布式	AWGR、TunableLaser	光纤	高	实验原型
Spanke	OCS	分布式	WSS	光纤	高	实验原型
ProjecToR	OCS	分布式	DMD	空间光	高	实验原型
FireFly	OCS	分布式	SM/GM	空间光	低	实验原型

资料来源：唐雄燕等《智算数据中心光电交换技术综述》、开源证券研究所

以谷歌 OCS 解决方案为例，OCS 在谷歌基础设施中主要有 Jupiter 数据中心和 TPU 数据中心两大应用场景。在初代 Jupiter 的基础上，通过引入 OCS 取代 Spine 层传统电交换机，将网络逻辑拓扑 CLOS 架构演进到 Aggregation 层的直接光互联，由于 OCS 采用光交换，对传输的速率无感，通过进一步引入 WDM 和环形器等技术可

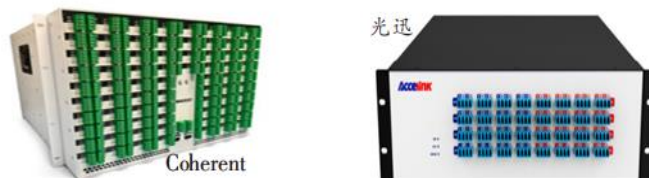
以实现在单根光纤上传输通道数的增加以及 Tx/Rx 双路信号，以提升单光纤的数据传输速率，在增加带宽容量的同时，减少电力消耗和降低成本。目前已有 Polatis、Coherent 和光迅等多家公司推出了商用的 OCS 产品。

图73：谷歌阿波罗架构核心层采用 OCS 交换机



资料来源：Mission Apollo: Landing Optical Circuit Switching at Datacenter Scale

图74：Coherent 和光迅科技推出商用 OCS 产品



资料来源：唐雄燕等《智算数据中心光电交换技术综述》

6、投资建议及交换机相关企业介绍

AIGC 有望拉动高速交换机需求持续增长，交换机产业链有望长期受益。

(1) 交换机&交换芯片推荐标的：紫光股份、盛科通信、中兴通讯；受益标的：锐捷网络、菲菱科思、共进股份、烽火通信、Arista 网络、思科、Juniper、博通、Marvell 等；

(2) 全光交换机受益标的：光迅科技、Coherent 等；

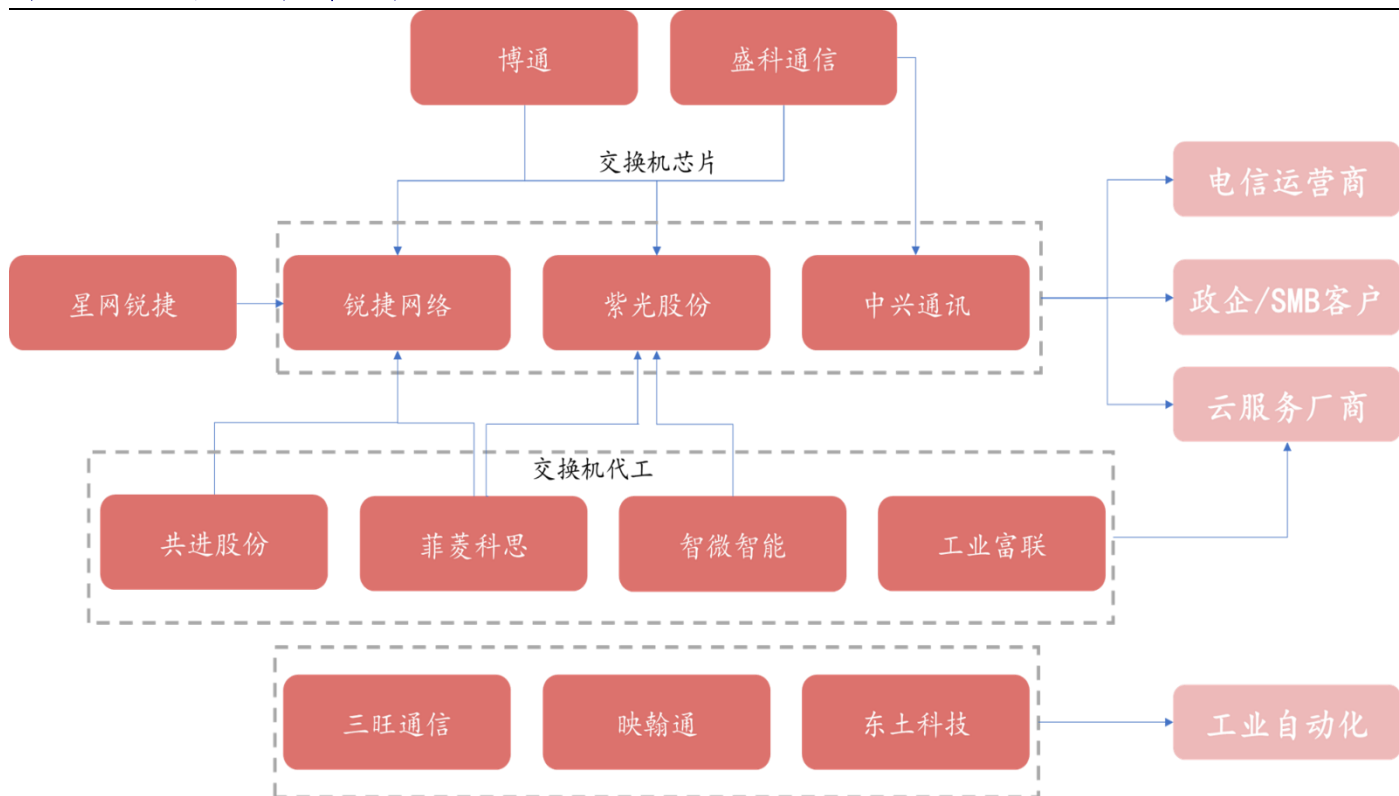
(3) 工业交换机受益标的：映翰通、三旺通信、东土科技等；

(4) 交换机配套 AIGC 推荐标的：宝信软件、润泽科技；受益标的：光环新网、奥飞数据、云赛智联、网宿科技等；

(5) 交换机配套光器件推荐标的：中际旭创、新易盛、天孚通信；受益标的：华工科技、光迅科技、源杰科技等；

(6) 交换机配套液冷推荐标的：英维克；受益标的：科华数据、网宿科技、飞荣达、高澜股份、申菱环境等。

图75：交换机行业重要上市公司



资料来源：开源证券研究所

表8：交换机行业受益标的估值表

证券简称	证券代码	评级	收盘价 (元)	市值 (亿元)	EPS (元/股)			PE		
					2024E	2025E	2026E	2024E	2025E	2026E
紫光股份	000938.SZ	买入	24.86	711.02	0.81	1.08	1.29	30.7	23.0	19.3
中兴通讯	000063.SZ	买入	31.06	1,251.11	1.89	2.07	2.35	16.4	15.0	13.2
锐捷网络	301165.SZ	未评级	48.37	274.83	0.95	1.24	1.52	51.1	39.1	31.8
烽火通信	600498.SH	增持	18.18	215.34	0.57	0.79	0.98	31.8	23.0	18.5
共进股份	603118.SH	增持	8.41	66.21	0.14	0.25	0.34	58.7	33.5	24.7
菲菱科思	301191.SZ	未评级	78.80	54.64	2.26	3.00	3.86	34.9	26.3	20.4
光迅科技	002281.SZ	增持	40.72	323.15	0.98	1.37	1.72	41.6	29.7	23.7
映翰通	688080.SH	未评级	35.61	26.30	1.54	1.95	2.46	23.1	18.3	14.5
三旺通信	688618.SH	未评级	23.75	26.21	1.06	1.46	1.95	22.4	16.3	12.2
东土科技	300353.SZ	未评级	13.66	83.99	0.08	0.15	0.25	162.8	91.1	55.3

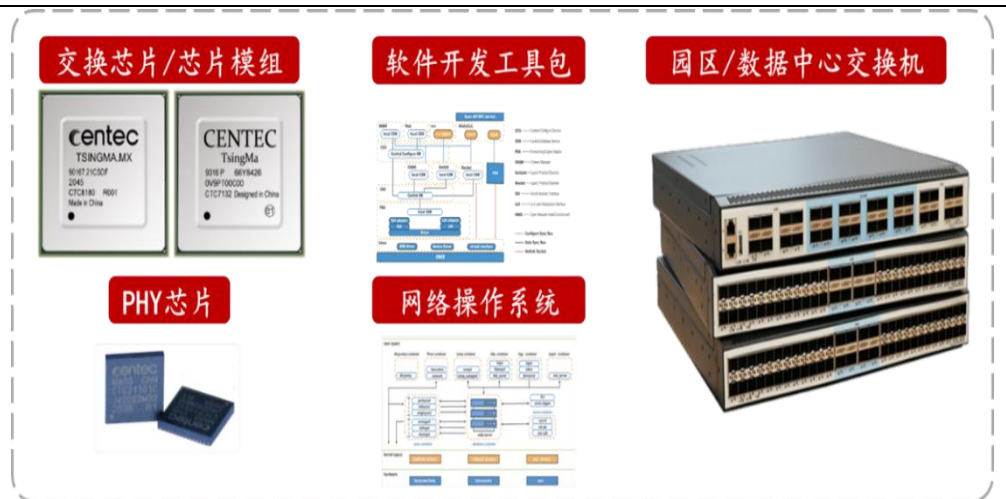
证券简称	证券代码	评级	收盘价 (元)	市值 (亿元)	营业收入 (亿元)			PS		
					2024E	2025E	2026E	2024E	2025E	2026E
盛科通信-U	688702.SH	买入	72.02	295.28	11.62	15.36	20.88	25.4	19.2	14.1

数据来源：Wind、开源证券研究所，股价为 2024 年 12 月 9 日收盘价（除紫光股份、盛科通信-U、中兴通讯为开源证券研究所预测外，其余均采用 wind 一致性预期）

6.1、盛科通信：稀缺的国产商用交换机芯片龙头

盛科通信立足于以太网交换芯片领域，提供以太网交换芯片模组和定制化产品解决方案。公司深耕以太网交换芯片领域的研发、设计和销售，现已形成丰富的以太网交换芯片产品序列，目前已拥有 100Gbps-2.4Tbps 的交换容量和 100M-400G 的端口速率的交换芯片，覆盖从接入层到核心层的以太网交换产品，自研以太网交换芯片已进入新华三、锐捷网络、迈普技术等国内主流网络设备商的供应链；在自研以太网芯片的基础上提供芯片模组，同时为行业客户提供定制化服务，提供定制化产品解决方案。此外，公司提供少量以白盒以太网交换机为主的终端产品，主要面向企业网络、运营商网络、数据中心网络和工业网络等场景需求。

图76：盛科通信产品包括交换芯片/模组、操作系统和交换机



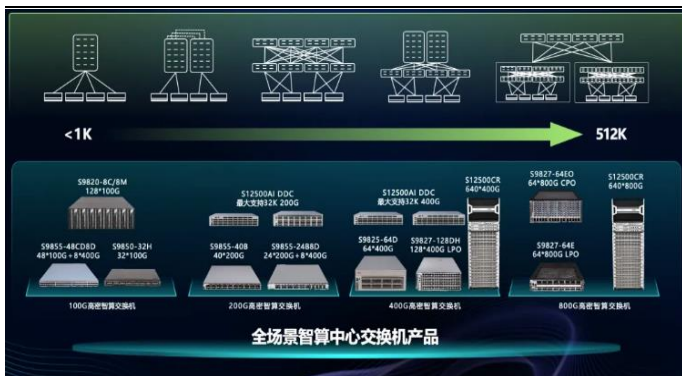
资料来源：盛科通信官网、开源证券研究所

6.2、紫光股份：国内交换机排头兵，率先布局 1.6T 端口、800G CPO 交换机

紫光股份作为云计算基础设施建设和行业智慧应用服务的领先者，提供智能化的网络、计算、存储、云计算、安全和智能终端等全栈 ICT 基础设施及服务，主要包含：（1）网络设备：交换机、路由器、WLAN、IoT、SDN、5G 小站、PON 等；（2）服务器：通用服务器、人工智能服务器、融合刀片服务器、关键业务服务器、边缘服务器等；（3）存储产品：全闪存集中式存储、混合闪存集中式存储、分布式存储、备份、容灾设备和存储网络设备等；（4）云计算与云服务：虚拟化平台、云操作系统、超融合产品、分布式存储系统、云桌面、大数据平台等产品及云运营服务等；（5）网络安全产品及服务：边界安全（包括防火墙、入侵防御系统等）、云安全、安全大数据、终端安全等领域产品及专业安全服务；（6）智能终端：商用笔记本电脑、商用台式机、智慧屏等。

交换机方面，公司深耕交换机产品多年，产品覆盖企业及园区、数据中心、工业交换机等场景，率先推出 51.2T 800G CPO 硅光数据中心交换机，首发 1.6T 端口智算交换机，适用于 AIGC 集群或数据中心高性能核心交换等业务场景。据 IDC 数据，2024Q1 公司在中国以太网交换机、企业网交换机、园区交换机市场，分别以 34.8%、36.5%、41.6% 的市场份额排名第一，在中国数据中心交换机市场份额 29.0%，位列第二。

图77：新华三数据中心交换机最大支持 51.2 万卡组网



资料来源：新华三公众号

图78：2023 年首发 800G CPO 交换机



资料来源：新华三官网

6.3、中兴通讯：核心芯片自研，400G 交换机全场景布局

中兴通讯拥有 ICT 行业完整的、端到端的产品和解决方案，处于行业领先地位。公司是全球领先的综合通信与信息技术解决方案提供商，基于 ICT 全栈核心能力，围绕连接（CT 技术）、算力（IT 技术）、云网融合构建高效数字底座，聚焦于运营商网络、政企业务和消费者业务。

运营商网络：基于 ICT 端到端全栈核心能力，面向运营商传统网络，推出 5G 基站、5G 核心网等无线产品；推出固网、光传输、路由器、交换机、光模块等有线产品；面向运营商云网业务，推出云电脑、服务器及存储产品、数据中心交换机和路由器、数据中心电力模块和液冷系统等产品。政企业务：紧跟算力浪潮，依托芯片、数据库和操作系统等底层核心技术，深入拓展国内政企市场，产品主要包括星云研发 AI 大模型、模型训推一体机、全系列服务器（AI 服务器、GPU 服务器、通用服务器、信创服务器等）及存储设备、数据中心交换机和路由器、数据库、5G+ 数字星云平台、车规级模组等。消费者业务：公司持续推出多款新品，产品矩阵不

断丰富，产品主要包含手机、平板电脑、AR 眼镜、MBB&FWA、PON CPE、机顶盒、路由器等。

交换机方面，公司推出国产超高密 400GE 框式交换机，搭载了自研的 7.2T 分布式转发芯片，采用业界领先的 112Gb/s 高速总线和正交连接器，整机支持高达 576 个 400G 或 288 个 800GE 接口，充分满足了大规模数据中心和云服务商对高带宽、高密度网络的需求。同时支持提供 51.2T/12.8T/3.2T/2T 等多种型号，灵活匹配不同场景下 100GE/400GE 组网需求。据 IDC 数据，2024Q1 公司在中国以太网交换机运营商市场营收实现同比增速第一，在数据中心交换机运营商市场领域，中兴通讯市场份额跃居第二位。

图79：中兴 400G 框式及盒式交换机



资料来源：中兴通讯官网

6.4、锐捷网络：中标多个头部互联网客户项目，受益于白盒化浪潮

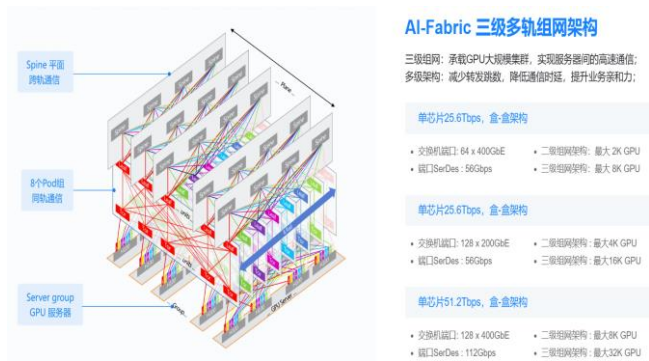
锐捷网络主营业务为网络设备、网络安全产品及云桌面解决方案的研发、设计和销售，主要产品包括网络设备（交换机、路由器、无线产品等）、网络安全产品（安全网关、下一代防火墙、安全态势感知平台等）、云桌面整体解决方案（云服务器、云终端、云桌面软件）以及 IT 运维等其他产品及解决方案。公司产品遍及网络建设中的各主要层级，广泛应用于局域网、城域网、广域网等各类型计算机网络中，根据 IDC 数据统计，2024 年第一季度，公司在中国以太网交换机市场占有率排名第三，在中国数据中心交换机市场占有率排名第三，在中国企业级 WLAN 市场占有率排名第二，其中 Wi-Fi 6 产品出货量排名第一。

图80：锐捷网络数据中心交换机产品矩阵



资料来源：锐捷网络官网

图81：锐捷网络推出 AI-Fabric 三级多轨网络架构



资料来源：锐捷网络官网

6.5、共进股份：800G 交换机陆续交付，突破多个海外客户

共进股份主要业务包括网通及数通业务（PON 系列、AP 系列、DSL 系列等各类宽带接入终端，交换机、核心路由等数通产品）、移动通信业务（4G/5G 小基站设备、固定无线接入设备以及以移动通信为技术基础的各类专业和综合应用产品）、汽车电子业务等。其中，数通业务已覆盖园区/SMB 交换机及 100G、400G、800G 等多种数据中心交换机。2024 年上半年公司 800G 数据中心交换机已开始陆续交付，工业交换机 JDM 项目完成零突破，获海外 P 客户、B 客户 SMB 交换机项目，交换机海外客户近 10 家。

图82：共进股份推出数据中心交换机产品



资料来源：共进电子官网

6.6、菲菱科思：发力中高端产品，国内领先 ODM/OEM 厂商

菲菱科思多年来专注于网络通信设备领域，在业务拓展方面，继续深耕园区接入、汇聚层中高端交换机，在产品方案方面，深度扩展国产方案交换机、安全防火墙、IOT 网关等业务；在中高端数据中心交换机产品部分，在 200G/400G /2.0T/8.0T 数据中心交换机上迭代 12.8T 等产品形态，扩展了基于国产 CPU 的 COME 模块；在交换机细分领域，扩展了工业控制和边缘计算场景需求的新一代 TSN 工业交换、Multi-GE (2.5G/5G/10G) 电口交换机及 2.5G 光上行千兆交换机/2.5G 光下行万兆上行全光交换机等，已经成为 S 客户、新华三、锐捷网络等国内主要品牌商的合格供应商。

图83：菲菱科思推出 12.8T 交换机



资料来源：菲菱科思官网

7、风险提示

（1）云计算需求不及预期

若云计算需求不及预期，将会影响到国内云巨头、电信运营商对于网络设备的采购。

（2）数字经济增长不及预期

若国内数字经济增速放缓，则会影响园区及企业、工业对网络设备采购，进而影响交换机、路由器等网络设备放量。

（3）AI 发展不及预期

若 AIGC 发展不及预期，将会影响 AI 后端组网需求，进而影响高速率数据中心交换机的采购量。