

NVIDIA GB200：重塑服务器/铜缆/液冷/HBM价值

——AI算力“卖水人”系列（3）

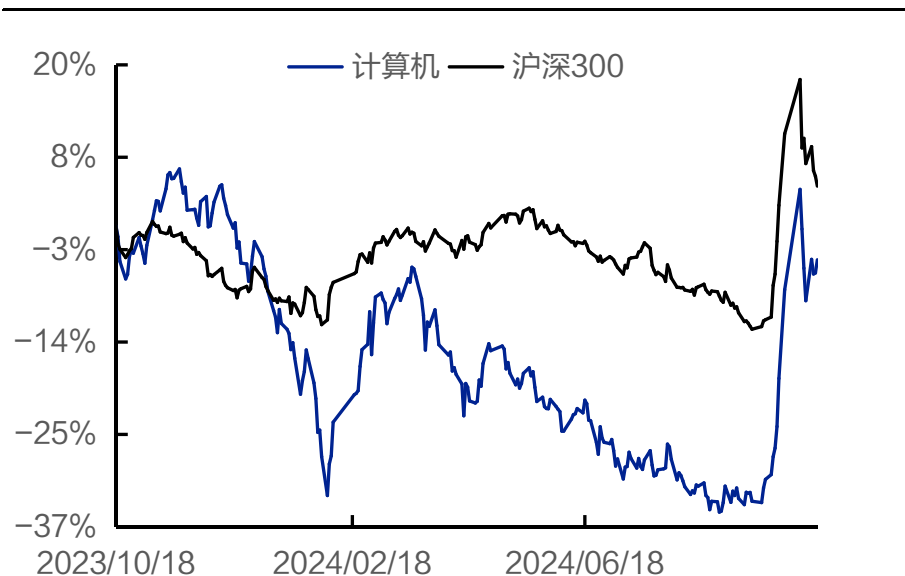
评级：推荐(维持)

刘熹(证券分析师)
S0350523040001
liux10@ghzq.com.cn

并购家 免费行业报告网

IPOIPO.CN 每日更新 不用注册会员 无广告 永久免费

最近一年走势



相对沪深300表现

表现	1M	3M	12M
计算机	44.1%	35.5%	-5.5%
沪深300	19.9%	8.2%	4.1%

相关报告

《计算机行业动态研究：鸿蒙原生：生态应用“裂变”生长，万物互联时代已来（推荐）*计算机*刘熹》——2024-09-24

《AI算力行业月度跟踪（202409）：OpenAI o1开创AI算力新纪元，Blackwell产能持续扩大（推荐）*计算机*刘熹》——2024-09-18

《计算机行业2024年中报总结：行业整体加速复苏，关注：算力、出海、华为链（推荐）*计算机*刘熹》——2024-09-11

- ◆ **核心要点：NV Blackwell系列或于2024Q4推出，有望带来机柜、铜缆、液冷、HBM四个市场的价值量提升。**
- ◆ **一、Blackwell系列：GB200计算能力远超H100，CSP厂商资本开支提升**
 - ✓ B200集成2080亿个晶体管，采用台积电N4P制程，为双芯片架构，192GB HBM3E，AI算力达20petaFLOPS (FP4)，是Hopper的5倍。GB200机架提供4种不同主要外形尺寸，每种尺寸均可定制。与H100相比，GB200 NVL72将训练速度(如 1.8 T 参数 GPT-MoE) 提高了 30 倍。
 - ✓ 互联网资本开支方面，Meta、Alphabet（谷歌）、微软、亚马逊等厂商均增加2024年资本开支指引。微软预计FY25Q1资本支出将环比增加，及2025财年资本支出将高于24财年；谷歌表示2024年季度资本开支将等于或高于一季度水平；Meta将2024年资本支出预测提高至370-400亿美元，亚马逊表示2024年下半年的资本投资将更高。
- ◆ **二、服务器细节拆分：主板从HGX到MGX，GB200 NVL72价值量提升**
 - ✓ GB200主板从HGX模式变为MGX，HGX是NVIDIA推出的高性能服务器，通常包含8个或4个GPU，MGX是一个开放模块化服务器设计规范和加速计算的设计，在Blackwell系列大范围使用。MGX模式下，GB200 Switch tray主要为工业富联生产，Compute Tray为纬创与工业富联共同生产，交付给英伟达。据Semianalysis，有望带来机柜集成、HBM、铜连接、液冷等四个市场价值量2-10倍提升。
- ◆ **三、铜连接：DACs市场较快增长，GB200 NVL72需求较大**
 - ✓ 高速线缆中，有源光缆适合远距离传输，直连电缆高速低功耗。据LightCounting，高速线缆规模预期28年达28亿美元，DACs保持较快增长，Nvidia的策略是尽可能多地部署DACs。据QYResearch，2022年，全球前十强厂商大约占据了 69.0% 的市场份额，其中安费诺的市场份额位居全球第一，国内厂商包括立讯精密、兆龙互联、金信诺、电联技术等，但全球市占率较低。
- ◆ **四、HBM：HBM3E将于下半年出货，英伟达为主要买家**
 - ✓ HBM 目前已经量产的共有HBM、HBM2、HBM2E、HBM3、HBM3E 五个子代标准。其中HBM3E将于下半年出货，HBM4或于2026年上市。预期2024年底HBM TSV产能约250K/m，三星与海力士扩产积极，三星HBM总产能至年底将达约130K（含TSV）；SK海力士约120K。英伟达目前是HBM最大买家，海力士与三星完成HBM3E验证。
- ◆ **五、冷板式液冷较为成熟，GB200 NVL72采用液冷方案**
 - ✓ AI大模型训推对芯片算力提出更高要求，提升单芯片功耗，英伟达B200功耗超1000W、接近风冷散热上限。液冷技术具备更高散热效率，包括冷板式与浸没式两类，其中冷板式为间接冷却，初始投资中等，运维成本较低，相对成熟，英伟达GB200 NVL72采用冷板式液冷解决方案。

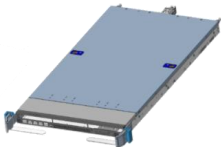
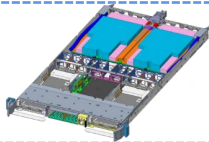
- 投资建议：大模型训推带动 AI算力需求增长，GB200等新一代算力架构将推出，算力产业链中的AI芯片、服务器整机、铜连接、HBM、液冷、光模块、IDC等环节有望持续受益。维持对计算机行业“推荐”评级。
- 相关公司
 - ✓ 1) **AI芯片**：海光信息、寒武纪、龙芯中科、景嘉微、英伟达、AMD、Intel
 - ✓ 2) **服务器整机**：工业富联、浪潮信息、中科曙光、华勤技术、中国长城、高新发展、神州数码、烽火通信、拓维信息、纬创、广达、英业达、纬颖、超微电脑。
 - ✓ 3) **服务器组件**：①散热：飞荣达、曙光数创、英维克、同方股份、申菱环境、高澜股份、奇鋆科技、双鸿、VERTIV；②主板：沪电股份、深南电路、胜宏科技、技嘉、华擎；③HBM：SK海力士、三星、美光、赛腾股份、联瑞新材；④铜连接：安费诺、沃尔核材、华丰科技。
 - ✓ 4) **光模块**：天孚通信、中际旭创、新易盛、光迅科技、华工科技。
 - ✓ 5) **数据中心**：奥飞数据、光环新网、宝信软件、数据港、电科数字。
- 风险提示：宏观经济影响下游需求、大模型产业发展不及预期、市场竞争加剧、中美博弈加剧、相关公司业绩不及预期等，各公司并不具备完全可比性，对标的相关资料和数据仅供参考。

投资建议与相关公司

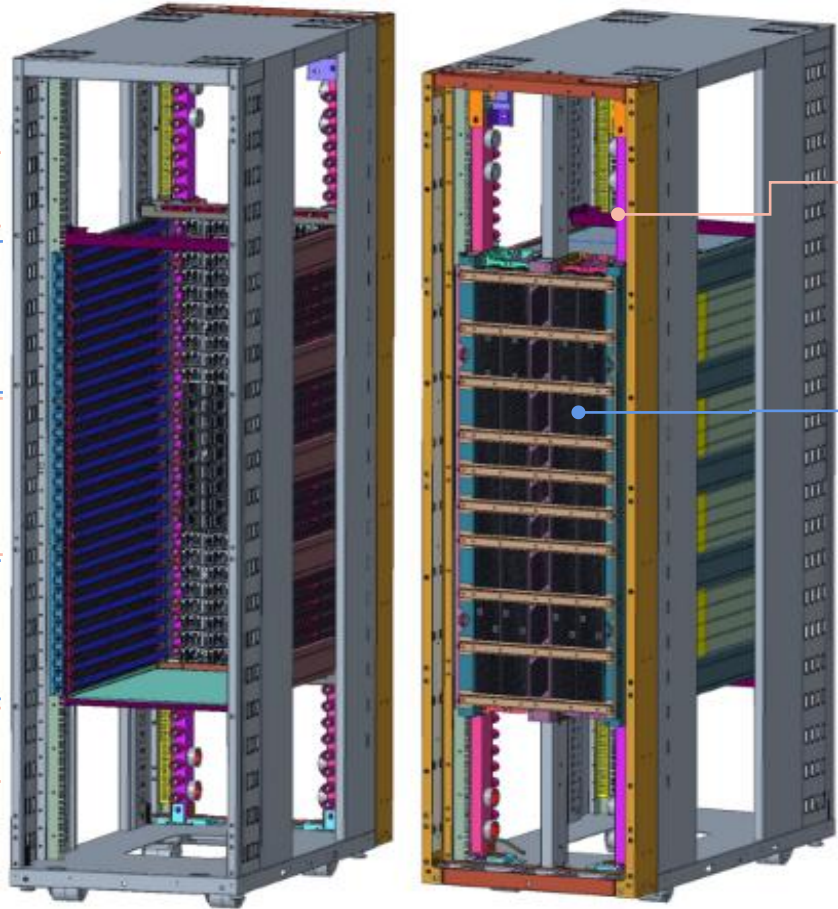


英伟达GB200 NVL72实物图

- **电源**
 - 内地：中国长城、欧陆通、杰华特、泰嘉股份
 - 中国台湾：台达
- **☆GPU**
 - 内地：华为海思、海光信息、寒武纪、龙芯中科。
 - 海外：英伟达、AMD、Intel
- **☆CPU**
 - 内地：华为海思、海光信息、飞腾、龙芯中科
 - 海外：英伟达、AMD、Intel
- **☆主板**
 - 内地：沪电股份、胜宏科技、深南电路、生益科技
 - 中国台湾：泰安电脑（神达）、技嘉、华擎、华硕、英业达、纬创、研华
 - 美国：Intel、Supermicro
- **Compute Tray * 10**
 - 内地：工业富联
 - 中国台湾：纬创
- **Switch Tray * 9**
 - 内地：工业富联
 - 中国台湾：纬创
- **Compute Tray * 8**
- **NV Switch Chip\网卡等**
 - 海外：英伟达、Mellanox
- **电源**



- **存储 NAND**
 - 公司：长江存储、兆易创新、佰维存储、朗科科技
 - 海外：三星、西部数据、铠侠、SK海力士、美光



正面

背面

- **光模块**
 - 内地：中际旭创、新易盛、光迅科技、天孚通信

- **☆分歧管（液冷：支持130KW的制冷能力）**
 - 内地：曙光数创、飞荣达、中航光电、英维克、同飞股份、申菱环境、高澜股份、依米康
 - 中国台湾：奇铤、双鸿、健策、建准、高力
 - 海外：VERTIV

- **☆线缆（铜连接：5000+根NVLink线缆）**
 - 内地：沃尔核材、华丰科技、兆龙互连、鼎通科技、立讯精密、神宇股份
 - 海外：安费诺

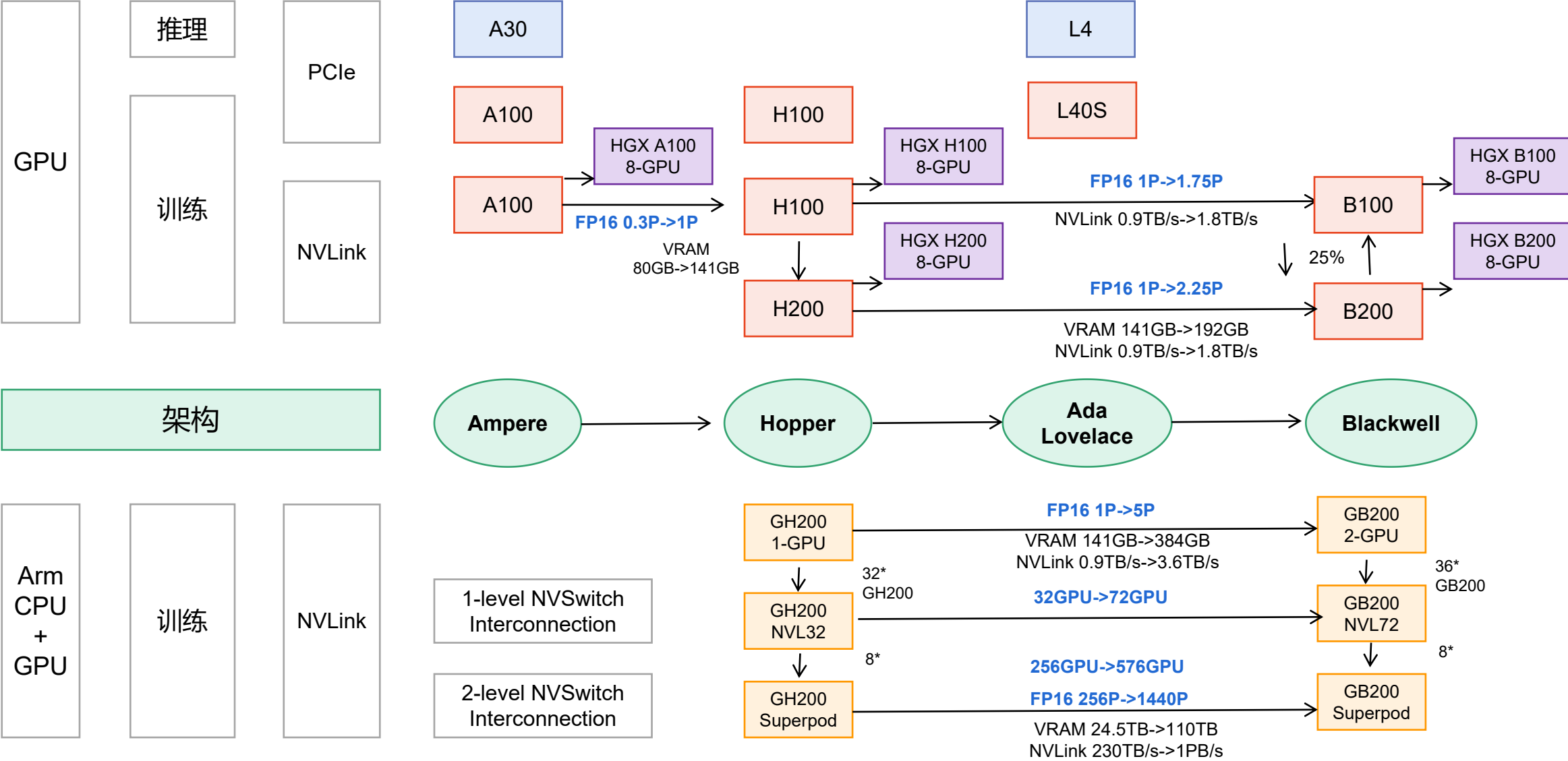
- **☆组装测试**
 - 整机厂商ODM：1) 内地：浪潮信息、工业富联、中科曙光、紫光股份、中国长城、华勤技术、神州数码、拓维信息、烽火通信、软通动力；2) 中国台湾：鸿海、广达、纬创、英业达、纬颖；3) 美股：戴尔、超微电脑
 - BIOS/BMC：1) 内地：卓易信息；2) 中国台湾：新唐、系微、信骅

注：蓝字为零部件。上图展示以英伟达GB200 NVL72整机柜为例，仅列示了部分参与公司

第一章 Blackwell 系列

GB200计算能力远超H100，CSP厂商资本开支提升

1.1 NVIDIA每一代GPU的计算能力、NVLink、内存持续扩大



1.1 NVIDIA每一代GPU的计算能力、NVLink、内存持续扩大

- 英伟达Blackwell GPU支持FP4精度，GB200的FP4计算能力可以达到20P，其计算能力是FP8的两倍，NVlink为3.6TB/s，显存容量为384GB，显存带宽为16TB/s。
- A100 -> H100: FP16密集计算能力增加了3倍以上，功耗从400W增加到700W。
- H200 -> B200: FP16密集计算能力增加了2倍以上，功耗从700W增加到1000W。
- B200的FP16密集计算能力约为A100的7倍，但功耗只增加了2.5倍。

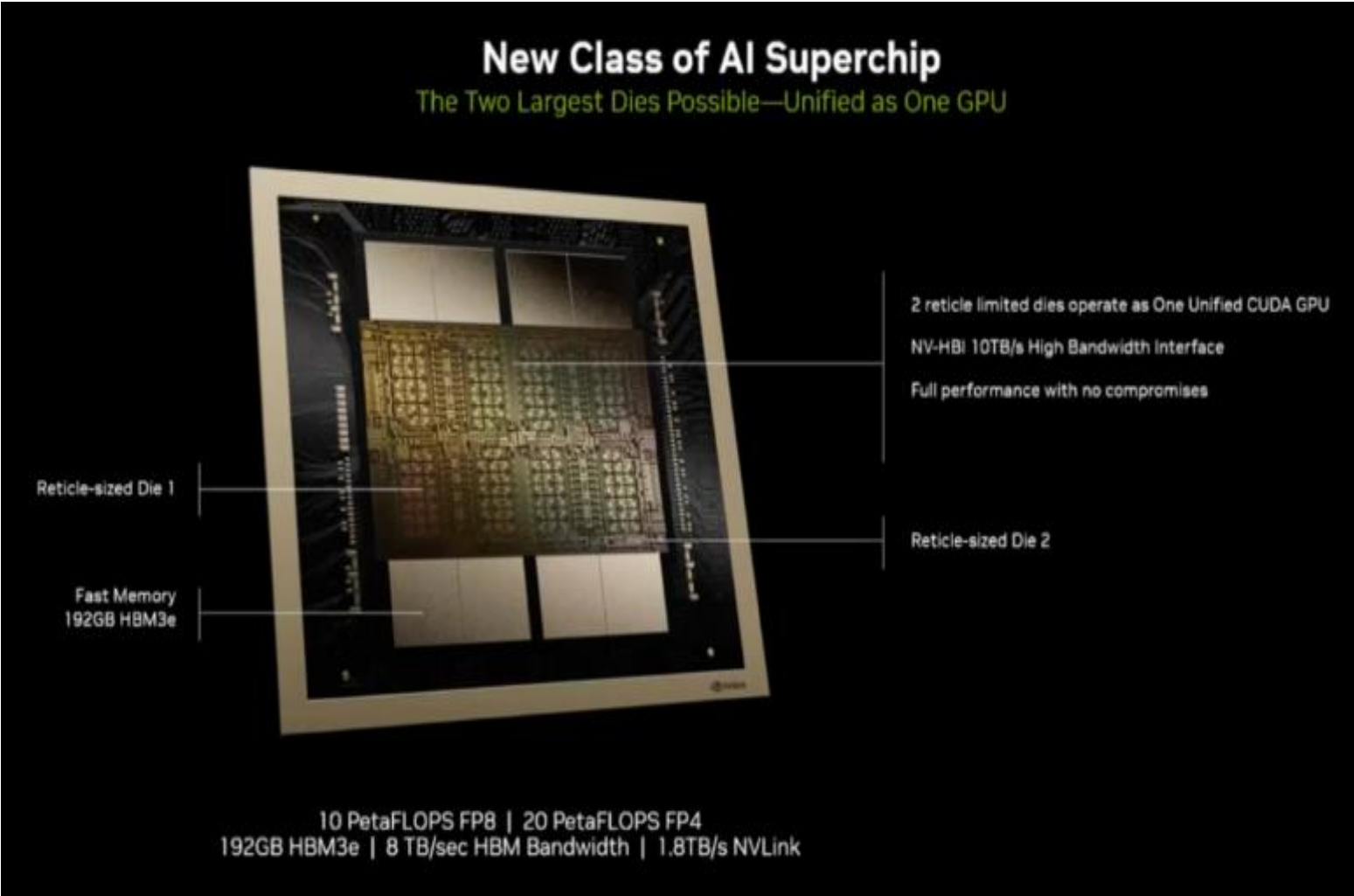
表：英伟达GPU迭代情况

Architecture	A100	H100	H200	GH200	B100	B200	Full B200	GB200
	Ampere	Hopper			Blackwell			
VRAM Size	80GB	80GB	141GB	96GB/144GB	180/192GB	180/192GB	192GB	384GB
VRAM Bandwidth	2TB/s	3.35TB/s	4.8TB/s	4TB/s/4.9TB/s	8TB/s	8TB/s	8TB/s	16TB/s
FP16(FLOPS)	312T	1P	1P	1P	1.75P	2.25P	2.5P	5P
INT8(OPS)	624T	2P	2P	2P	3.5P	4.5P	5P	10P
FP8(FLOPS)	X	2P	2P	2P	3.5P	4.5P	5P	10P
FP6(FLOPS)	X	X	X	X	3.5P	4.5P	5P	10P
FP4(FLOPS)	X	X	X	X	7P	9P	10P	20P
NVLink Bandwidth	600GB/s	900GB/s	900GB/s	900GB/s	1.8TB/s	1.8TB/s	1.8TB/s	3.6TB/s
Power Consumption	400w	700w	700w	1000w	700w	1000w	1200w	2700w
Notes	Die*1	Die*1	Die*1	Grace CPU*1 H200 GPU*1	Die*2	Die*2	Die*2	Grace CPU*1 Blackwell GPU*2

1.1 B200：Blackwell为TSMC 4N工艺，B200采用双芯片封装

表：英伟达B200 GPU

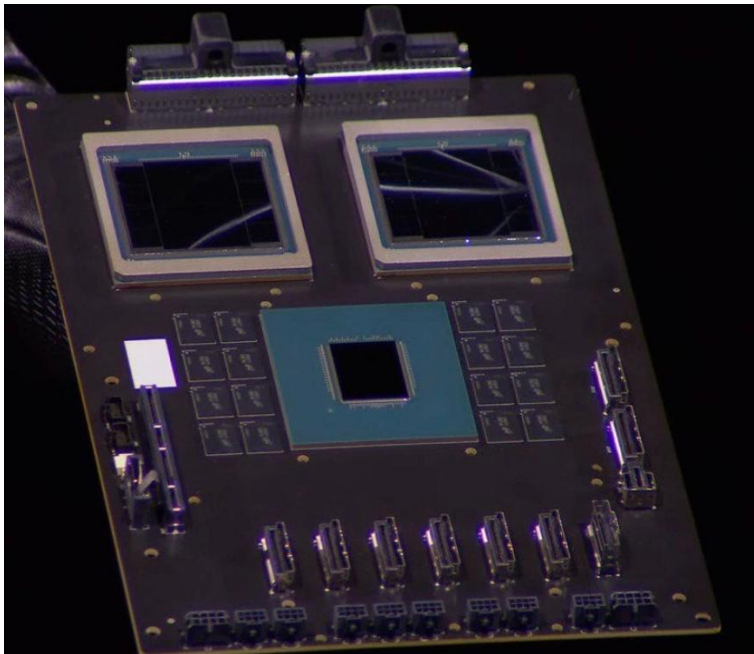
- B200 GPU：
- **工艺**：Blackwell GPU采用TSMC的N4P技术，H100 GPU采用N4工艺。H100是一个单芯片（单个完整的半导体单元）封装，Blackwell GPU是一个多芯片封装，有2个芯片。
- **计算能力**：每个Blackwell GPU芯片的FP8计算能力大约是H100的2.5倍。
- **通信能力**：B200为双芯片架构，两个芯片之间的通信带宽为10TB/s。连接8个8层堆叠的HBM3E，容量达到了192GB。



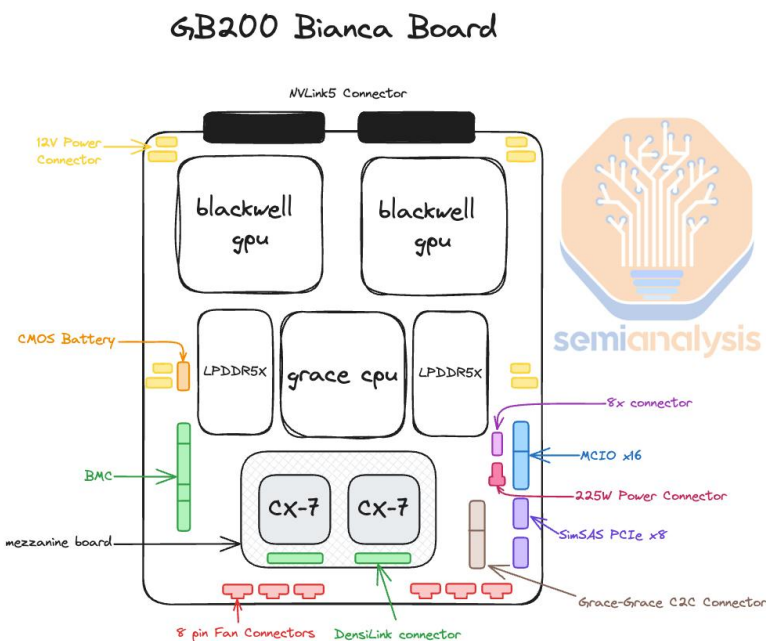
1.1 GB200：计算性能可达H100 30倍，NVLink-C2C互联900GB/s

- **GB200性能**：相较于H100，GB200的算力提升了6倍，在处理多模态特定领域任务时，其算力更是能达到H100的30倍。GB200集成了诸多先进技术，包括第二代 Transformer引擎、第五代 NVLink 高速互联技术等。
- **GB200组成**：是将2个B200芯片和1个GraceCPU整合到一起，相应的 GPU 算力和显存都加倍。CPU 和 GPU 之间依然通过 900GB/s 的 NVLink-C2C 实现高速互联，对应的功耗为 2700W。

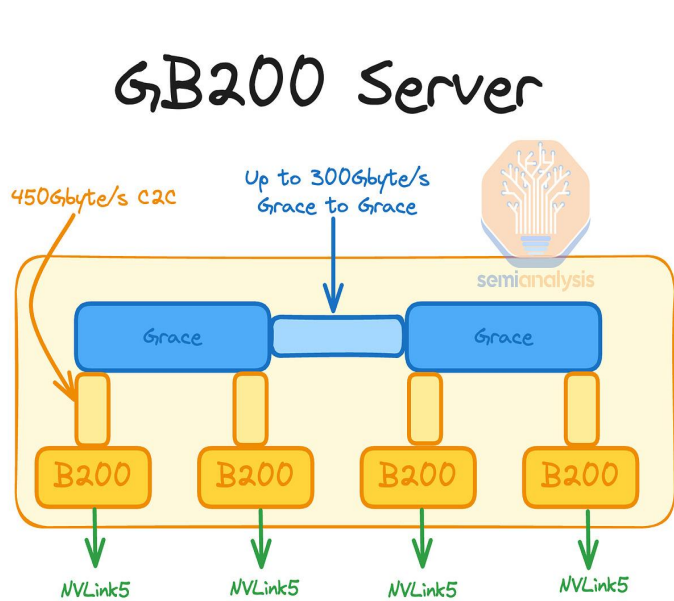
图：英伟达GB200实物图



图：英伟达GB200结构图



图：英伟达GB200 Compute Tray结构图



1.2 Rubin将于2026年推出，GPU产品迭代加速

- 2024年6月2日，COMPUTEX 2024大会召开，英伟达CEO黄仁勋在大会上展示GPU、CPU、NVLink等方面技术路线图。
- 1) GPU：一年迭代一次。下一代Rubin于2026年上市：将集成8颗HBM4；2027年推出Rubin Ultra GPU，将集成12颗HBM4。
- 2) CPU：英伟达确认下一代CPU平台Vera CPU将于2026年推出。
- 3) NVLink：第五代NVIDIA NVLink互连可扩展至576个GPU。2026年，计划推出第六代NVLink switch，达到3600GB/sec。
- 4) NVIDIA Spectrum™-X以太网网络平台：全球首款专为AI打造的以太网网络平台，计划每年都推出新的 Spectrum-X 产品。

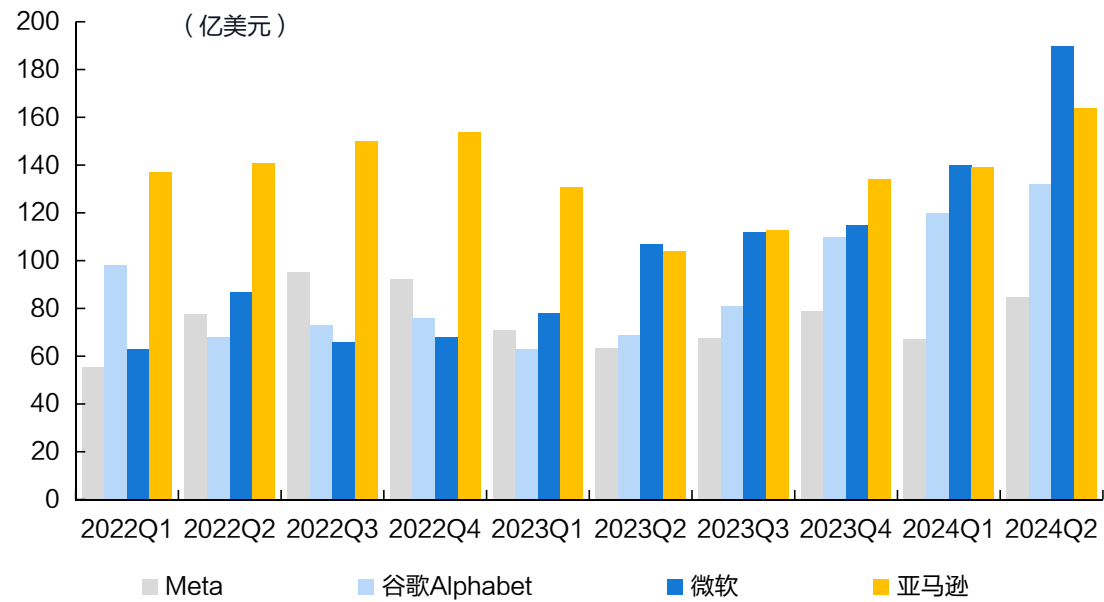
图：英伟达未来产品迭代情况



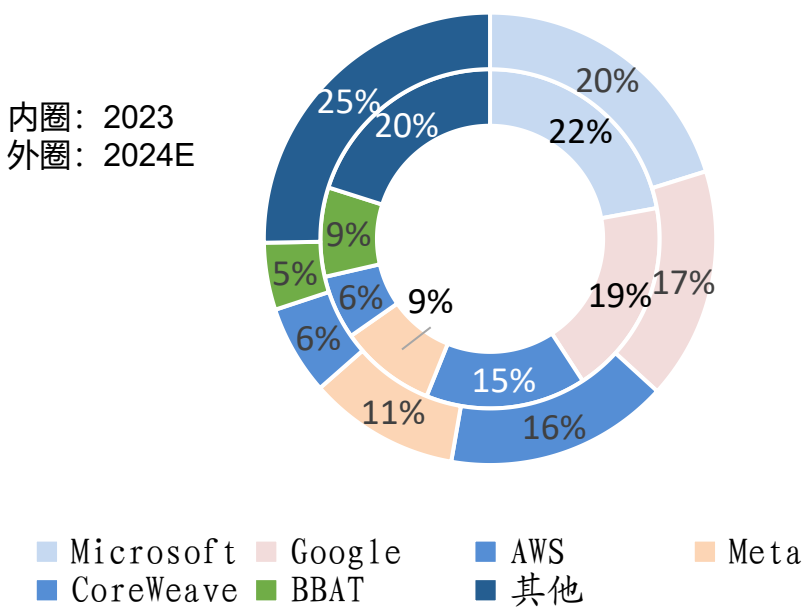
1.3 互联网厂商资本开支指引提升，ASIC AI服务器采购占比增长

公司	2024年互联网大厂资本开支预期情况
Microsoft	2024Q2资本开支为190亿美元。资本开支预期大约一半用于基础设施需求，建设和租赁数据中心，其余相关支出主要用于服务器，包括CPU和GPU。未来继续扩大基础设施投资，预计2024Q3资本支出将环比增加。
Alphabet（谷歌）	2024Q2公司资本开支达到130亿美元。公司预计全年每季度资本支出将大致维持第一季度120亿美元或略高。
Meta	2024Q2公司资本开支为85亿美元。公司上调全年资本支出下限至370亿-400亿美元（此前预期为350亿-400亿美元），并预计2025年的资本支出将明显增加，基础设施是增加的主要因素。
亚马逊	2024Q2资本开支为164亿美元。公司预计下半年的资本投资将更高，大部分支出将用于支持对AWS 基础设施日益增长的需求

图：2022-2024Q2 各厂商资本性开支



图：2023-2024年全球CSP对高阶AI服务器需求占比



1.3 互联网厂商资本开支指引提升，ASIC AI服务器采购占比增长



- **CSP资本开支持续投向AI服务器采购。**据TrendForce预估，全球AI服务器2024年第2季出货量将季增近20%，全年出货量上修至167万台，年增率达41.5%。2024年大型CSPs及品牌客户等对于高阶AI服务器的需求较好，来自北美CSPs业者（如AWS、Meta等）持续扩大自研ASIC（专用集成电路），以及中国本土业者如：阿里巴巴、百度、华为等积极扩大自主ASIC 方案，促ASIC服务器占整体AI服务器的比重在2024年将提升至26%，而主流搭载GPU的AI服务器占比则约71%。

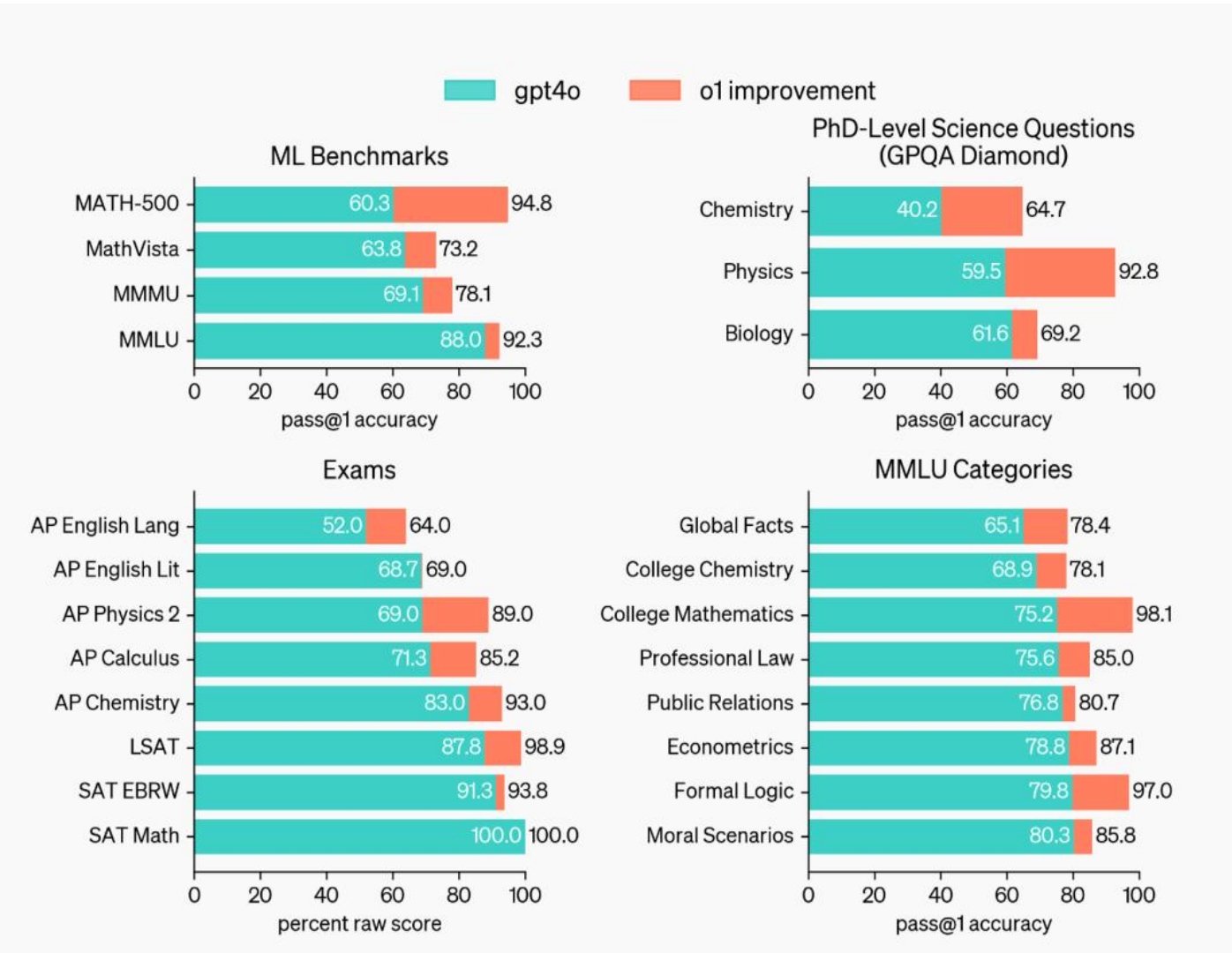
表1：2024年搭载ASIC芯片AI服务器出货占比将逾2.5成

公司	2022	2023	2024E
NVIDIA	67.6%	65.5%	63.6%
AMD（包括Xilinx）	5.7%	7.3%	8.1%
Intel（包括Altera）	3.1%	3.0%	2.9%
Others	23.6%	24.1%	25.3%
全部	100%	100%	100%

1.3 OpenAI模型加速迭代，o1推理能力提升算力需求

- OpenAI-o1拥有进化的推理能力。**9月12日，OpenAI发布“草莓”模型的部分预览版——OpenAI-o1，这是OpenAI计划推出的一系列“推理”模型中的第一个。OpenAI-o1可以进行通用复杂推理，擅长解决复杂问题，尤其是编码、数学、科学远超OpenAI-4o。
- OpenAI-o1和o1-mini模型已经在ChatGPT中上线，Plus和Team订阅用户可以直接体验。**开发者对o1的访问非常昂贵，在API中，o1-preview的价格是每100万个输入tokens 15美元，每100万个输出tokens 60美元。相比之下，GPT-4o的价格是每100万个输入tokens 5美元，每100万个输出tokens 15 美元。但目前的请求频率有限制，o1-preview的每周速率限制为30条消息，o1-mini的每周速率限制为50条。
- 我们认为，OpenAI-o1及o1 mini模型在编码、数学、科学等复杂推理任务方面的能力提升，将促进其提升应用效果、扩大应用范围，促进算力需求持续增长。**同时，OpenAI-o1在算法上的优化，以及在思维链推理、多步推理等能力的提升，或将导致OpenAI-o1对比上一代模型在同一任务下消耗更大的算力，更高的定价亦或反应了该趋势。因此，OpenAI-o1的推广有望全面扩大AI算力需求。

图：GPT-4o与o1的多方面基准测试对比



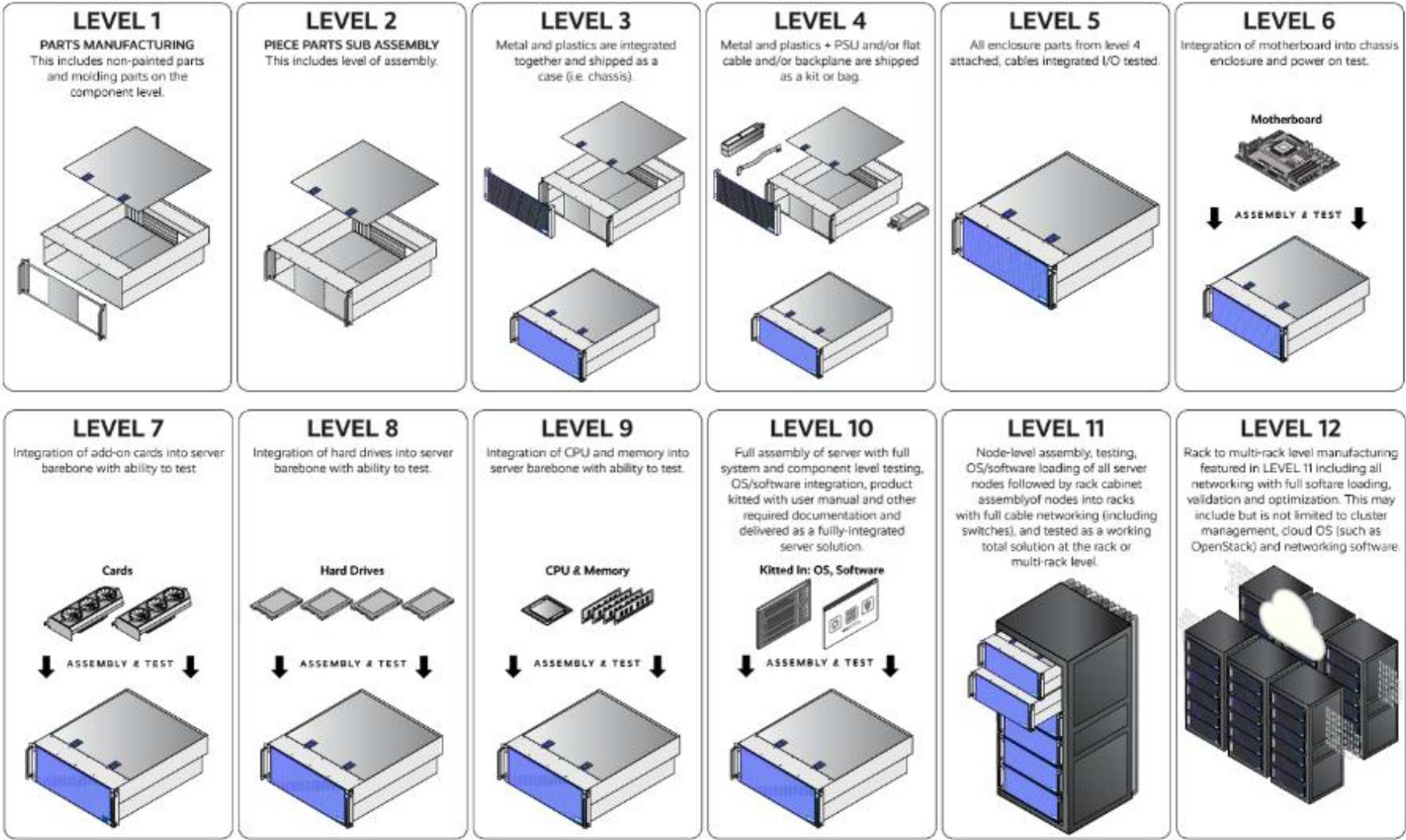
第二章 服务器细节拆分

主板从HGX到MGX，GB200 NVL72价值量提升

2.1 英伟达整机：HGX到MGX计算平台，服务器整机到系统形态

服务器制造级别可分为 Level 1-Level 12。
 Level 1为零部件制造，Level 6为主板集成，通常为 ODM 发货“服务器准系统”时提供的。
 Level 10为完整服务器组装和测试，能够达到 10 级制造水平的制造商将提供有效的服务器解决方案。Level 11-12级可将多台服务器联网作为机架级甚至多机架级解决方案。

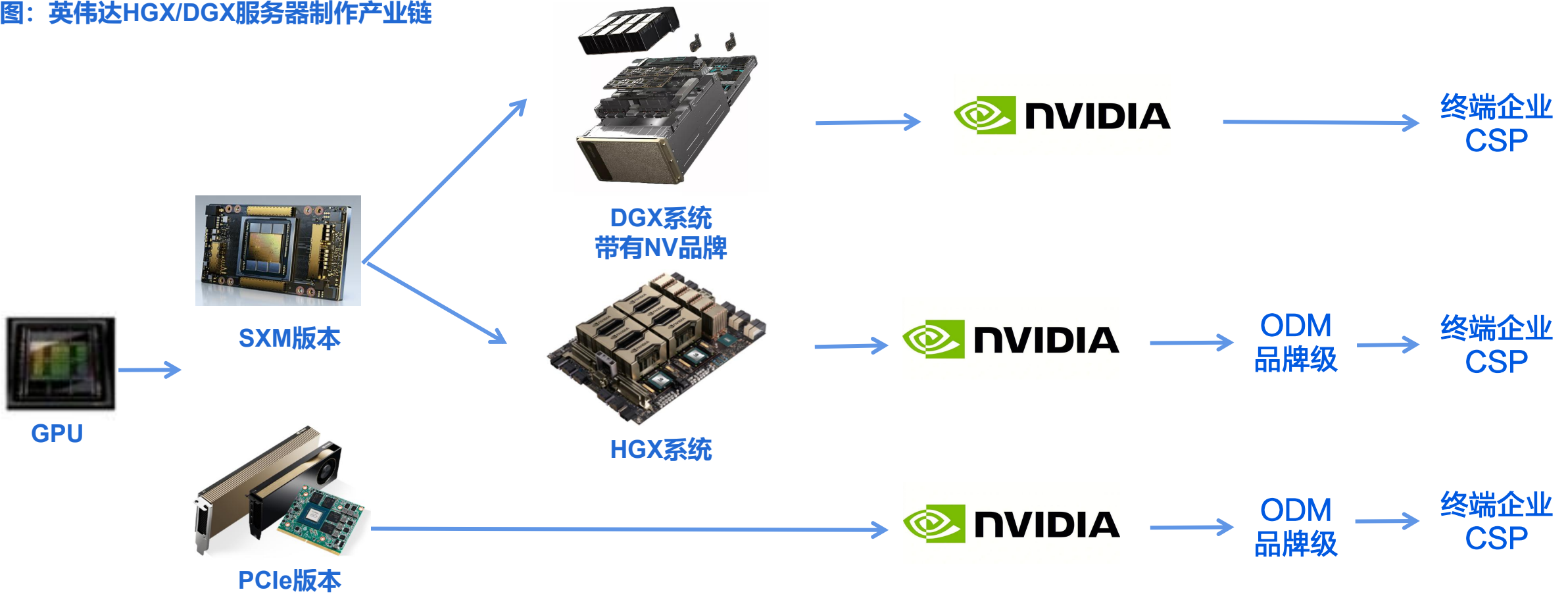
图：服务器制造等级分类



2.1 英伟达整机：HGX版本由ODM出货，DGX版本由英伟达交付

- **HGX**：模组主要为SXM版本，随后8个GPU SXM模组构建成一个HGX UBB基板，基本交付给英伟达后，分配给品牌级服务器厂商，包括鸿海、广达、纬创、英业达等厂商。
- **DGX**：SXM模组生产后，交给下游服务器厂商进行组装，整机交付给英伟达，流向下游CSP厂商与企业客户等。

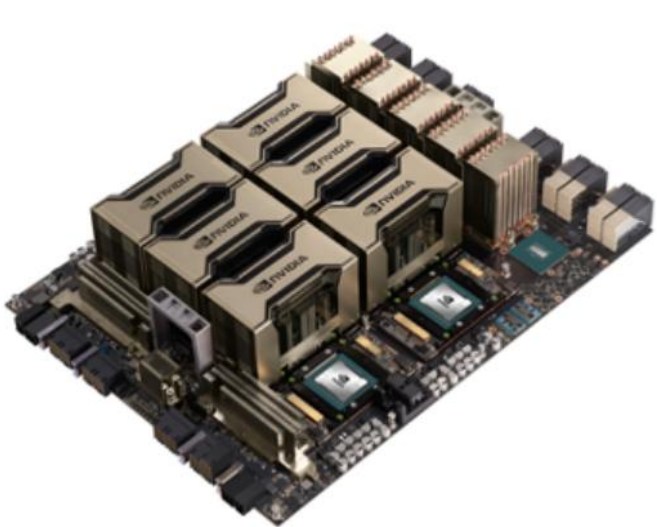
图：英伟达HGX/DGX服务器制作产业链



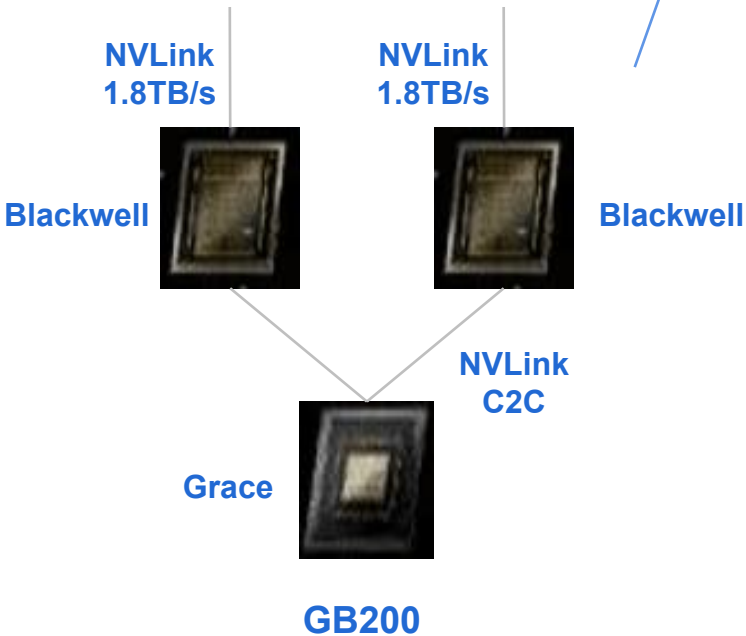
2.1 英伟达整机：MDX开放模块化设计，NVL72或采用此结构

- **MGX (Module GPU Accelerator)**：是一个开放模块化服务器设计规范和加速计算的设计。MGX让系统制造商快速且经济高效地为 AI、HPC 和 NVIDIA Omniverse 应用程序构建 100 多种服务器变体。**GB200 NVL72**整个系统由18个Compute Trays计算节点和9个Switch Trays交换节点组成。

图：英伟达服务器制作产业链从看HGX-8卡到MGX



HGX系统



GB200 SUPERCHIP
COMPUTE TRAY

2x GB200
80 PETAFLIPS FP4 AI INFERENCE
40 PETAFLIPS FP8 AI TRAINING
1728 GB FAST MEMORY
1U Liquid cooled
18 Per Rack



NVLINK SWITCH TRAY

2x NVLINK SWITCH
14.4 TB/s Total Bandwidth
SHARPv4 FP64/32/16/8
1U Liquid cooled
9 Per Rack



2.1 英伟达整机：Hopper到Blackwell，服务器系统多元变化

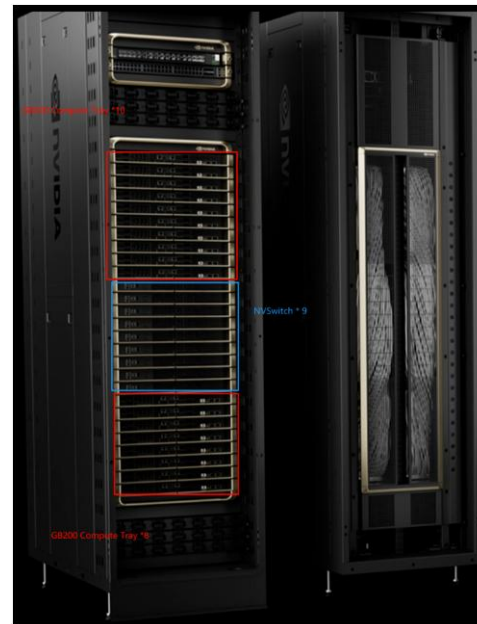
- 从Hopper平台升级到Blackwell平台，英伟达系统级解决方案存在变化：

- ✓ **DGX H100服务器**：采用8个H100 GPU，向外扩展18个 NVLink，总共900GB/s双向带宽；还包括 CPU主板、系统内存、网卡、PCIe交换机等组件。
- ✓ **GB200 NVL72机架**：NVL72包含18块计算托盘与9块交换托盘，每块计算托盘包括2个Grace CPU与4个B200 GPU，机柜内部通信主要采用NVLink网络架构，使用高速铜缆通信互连，提供水冷散热。

图：英伟达DGX H100服务器



图：英伟达GB200 NVL72整机柜



图：英伟达DGX Superpod



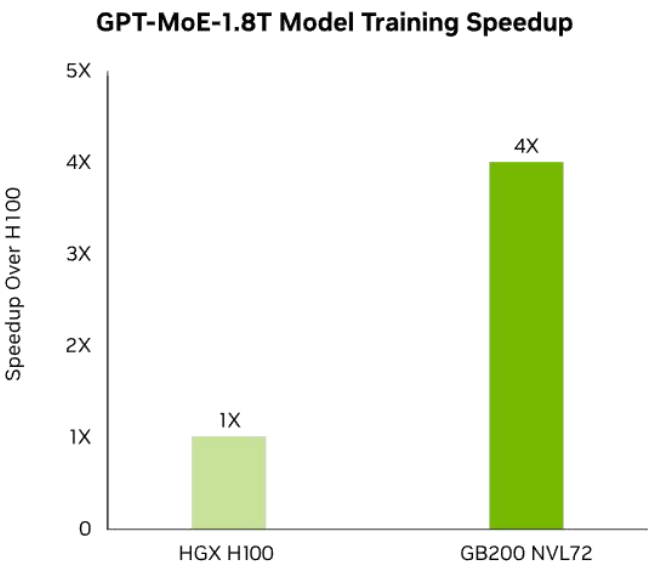
图：GB200 NVL72 Superpod



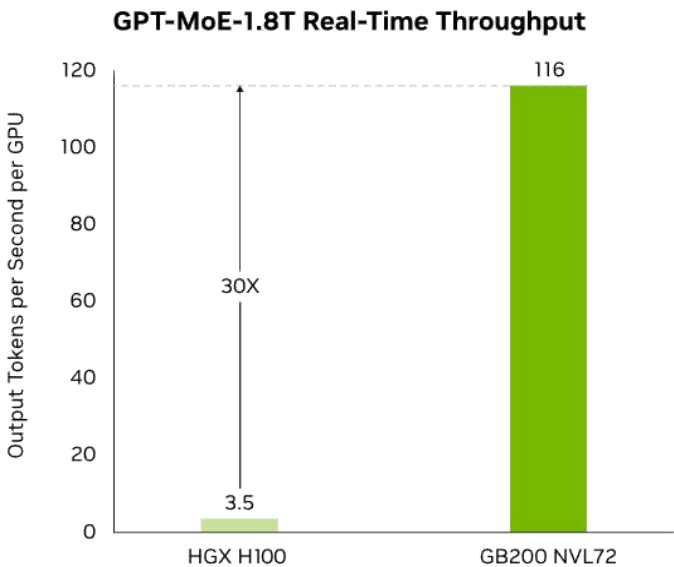
2.1 英伟达整机：GB200 NVL72 训练与推理能力倍数提升

- AI 训练：**GB200 包含速度更快的第二代 Transformer 引擎。与相同数量的 NVIDIA H100 GPU 相比，GB200 NVL72 可为 GPT-MoE-1.8 T 等大型语言模型提供 4 倍的训练性能。
- AI 推理：**GB200 引入了先进的功能和第二代 Transformer 引擎，可加速 LLM 推理工作负载。与上一代 H100 相比，它将资源密集型应用程序 (例如 1.8 T 参数 GPT-MoE) 的速度提高了 30 倍。新一代 Tensor Core 引入了 FP4 精度和第五代 NVLink 带来的诸多优势，使这一进步成为可能。
- 加快数据库查询速度：**GB200利用NVIDIA Blackwell 架构中具有高频宽记忆体效能的NVLink-C2C和专用解压缩引擎，将关键资料库查询的速度提升为CPU的18倍，总拥有成本降低5倍。

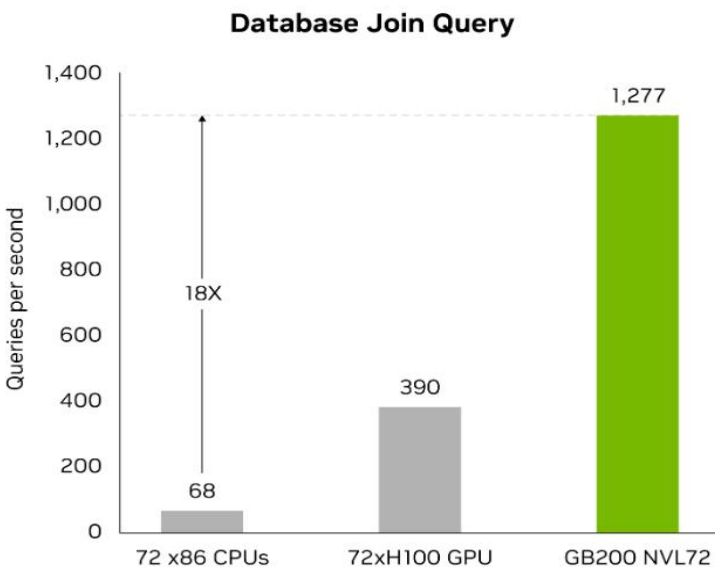
图：GB200 NVL72训练性能约为H100的4倍



图：对比H100，GB200可提供30倍实时吞吐量



图：GB200可提升数据库查询速度

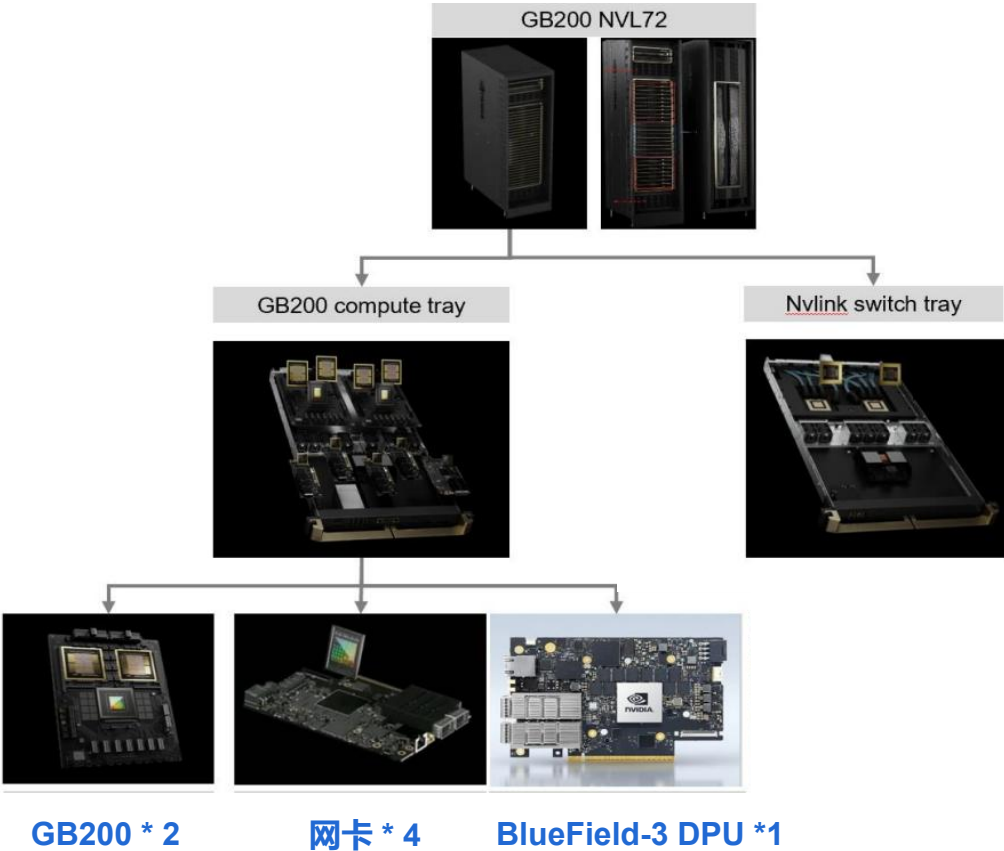


2.2 GB200 NVL系统：发展NVL72为代表的四种不同外形尺寸

- GB200机架提供4种不同主要外形尺寸，每种尺寸均可定制。
- ①GB200 NVL72：除了一家计划将其部署为主要变体的超大规模提供商外，据SemiAnalysis，在Blackwell Ultra之前很少部署此版本，因为即使使用直接到芯片的液体冷却（DLC）也是如此，大多数数据中心基础设施也无法支持如此高的机架密度。

- NVL72：18个1U计算托盘和9个NVSwitch托盘。计算托盘：2个 Bianca板，每个板是1个Grace CPU和2个Blackwell GPU。NVSwitch托盘：有两个28.8 Tb/s NVSwitch5 ASIC。120kW功率。

图：英伟达GB200 NVL72 的计算与网络节点



GB200 NVL72 实物图



GB200 NVL72 架构图



2.2 GB200 NVL系统：NVL36 x 2增加交换托盘与功耗

GB200 NVL36 x 2

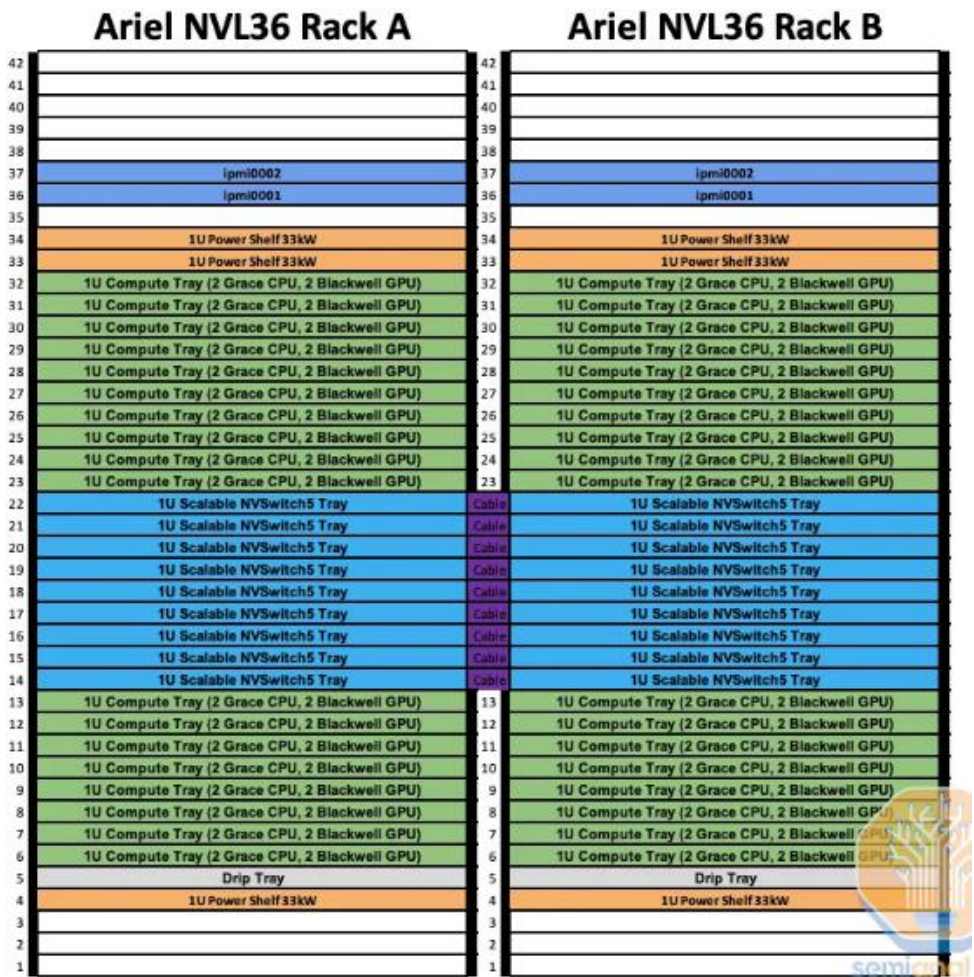
- ②GB200 NVL36x2：**每个机架18个Grace CPU 和36个Blackwell GPU
- 计算托盘：**每个计算托盘高 2U，包括2 个 Bianca 板
- NV Switch托盘：**两个28.8Tb/s NVSwitch5 ASIC 芯片。
 每芯片背板后方 14.4Tb/s，前板14.4Tb/s。每个NVswitch托架都有18个1.6T双端口OSFP壳体，水平连接到一对NVL36机架。
- 每机架66kW，总共为132kW。**由于额外的 NVSwitch ASIC 和跨机架互连布线的要求，与 NVL72 相比，NVL36x2 系统确实多使用了约10kW 的功率。NVL36x2 总共将有 36 个 NVSwitch5 ASIC，而 NVL72 上只有 18 个 NVSwitch5 ASIC。



2.2 GB200 NVL系统：推出NVL36x2（Ariel）与x86 B200版本

- ③ GB200 NVL36x2（Ariel）：**据 SemiAnalysis，主要由 Meta 使用。Meta 大部分采购普通的 NVL36x2，因为更适合 GenAI 工作负载，Ariel 版本将仅用于最大的推荐系统工作负载。由于 Meta 的推荐系统训练和推理工作负载，它们需要更高的 CPU 内核和更多的每 GPU 内存比率。
- ④ x86 B200 NVL72/NVL36x2：**据 SemiAnalysis，在 2025 年第二季度，将推出 B200 NVL72 和 NVL36x2 外形尺寸，它们将使用 x86 CPU，而不是 Nvidia 的内部 grace CPU。这种外形规格称为 Miranda。据 SemiAnalysis，每个计算托盘的 CPU 到 GPU 将保持不变，每个计算托盘 2 个 CPU 和 4 个 GPU。

GB200 NVL36x2（Ariel）



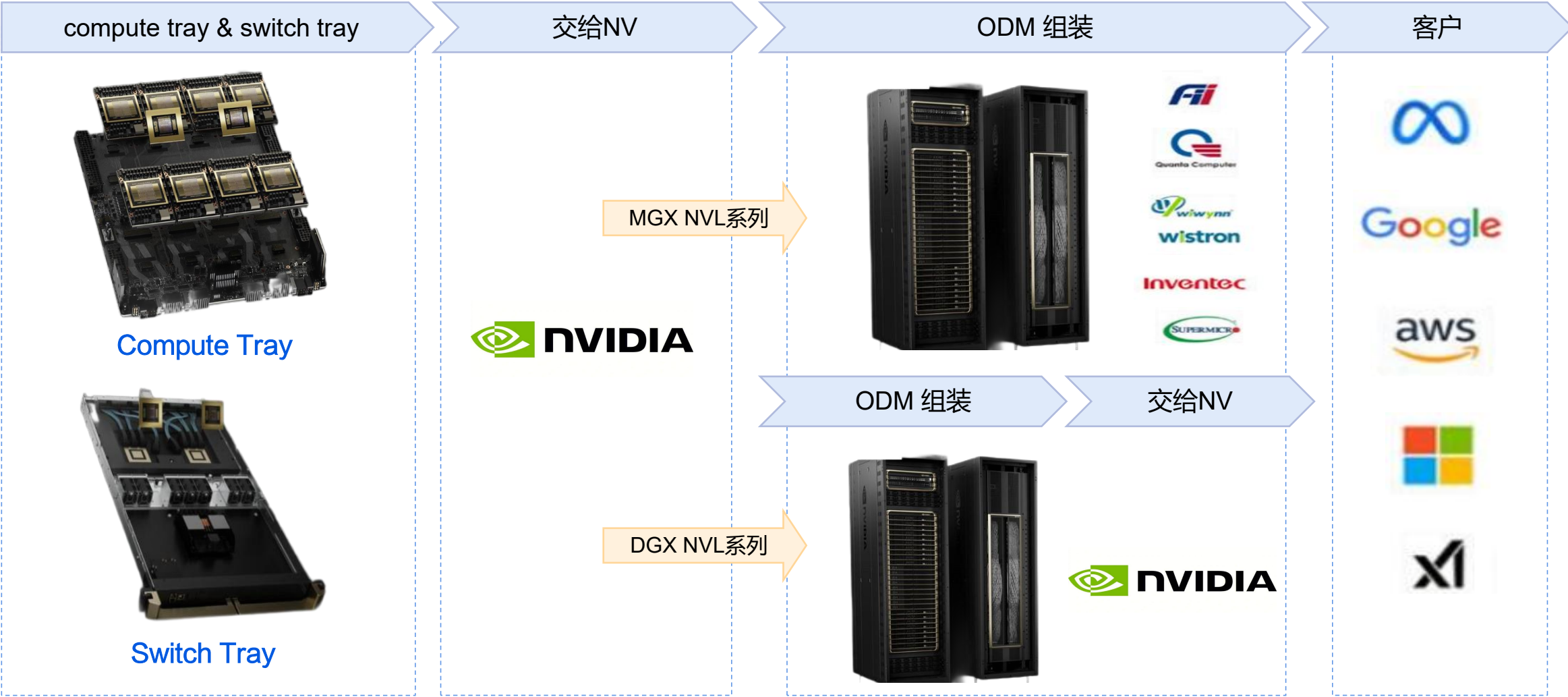
NVL36 x 2（Ariel）：
 每个机架 18 个 Grace CPU 和 36 个 Blackwell GPU

定制“Ariel”板：1 个 Grace CPU 和 1 个 Blackwell GPU

NVSwitch 托盘：有 18 个 1.6T 双端口 OSFP 笼，水平连接一对 NVL36 机架

2.2 NVL72组装模式：MGX-->整机柜-->集群

图：NVL72的组装流程



2.2 NVL72组装模式：GPU数量、计算能力、网络性能实现提升

● NVL系统

- ✓ NVL32 -> NVL72：GPU数量从32增加到72，FP16密集计算能力从32P增加到180P，几乎增加了6倍，而功耗也从40kW（估计数据）增加到120kW，近3倍。
- ✓ GH200 SuperPod -> GB200 SuperPod：GPU数量从256增加到576，FP16密集计算能力从256P增加到1440P，几乎增加了6倍。
- ✓ 在NVL72和GB200 SuperPod中使用最新的800Gb/s带宽ConnectX-8 IB网络卡，而HGX B100和HGX B200仍然使用400Gb/s带宽ConnectX-7 IB网络卡。
- ✓ 需要注意的是：NVIDIA介绍说GB200 SuperPod由8个NVL72组成，而GH200 SuperPod不是由8个NVL32组成。

表：Hopper架构和Blackwell架构对应的NVL机柜和SuperPod

Architecture	NVL32	GH200 SuperPod	NVL72	GB200 SuperPod
	32 xGH200	256 x GH200	36 x GB200	288 x GB200
	32 x GPU	256 x GPU	72 x GPU	576 x GPU
	Hopper		Blackwell	
HBM VRAM Size	32 x 144GB=4.6T	256 x 96GB=24.5TB	36 x 384GB=13.8TB	288 x 384GB=110TB
LPDDR5X VRAM Size	32 x 480GB=15.4T	256 x 480GB=123TB	36 x 480GB=17.3TB	288 x 480GB=138TB
HBM VRAM Bandwidth	32 x 4.9TB/s	256 x 4TB/s	72 x 8TB/s	576 x 8TB/s
FP16(FLOPS)	32P	256P	180P	1440P
INT8(OPS)	64P	64P	360P	2880P
FP8(FLOPS)	64P	64P	360P	2880P
FP6(FLOPS)	X	X	360P	2880P
FP4(FLOPS)	X	X	720P	5760P
NVSwitch System	Gen3-64 Port/NVSwitch		Gen4-72 Port/NVSwitch	
	9 x L1 NVSwitch Tray	96 x L1 NVSwitch Tray 36 x L2 NVSwitch Tray	9 x L1 NVSwitch Tray	144 x L1 NVSwitch Tray 72 x L2 NVSwitch Tray
GPU-to-GPU Bandwidth	0.9TB/s	0.9TB/s	1.8TB/s	1.8TB/s
NVLink Bandwidth	36 x 0.9T/s=32TB/s	256 x 0.9T/s=230TB/s	72 x 1.8T/s=130TB/s	576 x 1.8TB/s=1PB/s
Ethernet	16 x 200Gb/s	256 x 200Gb/s	18 x 400Gb/s	
InfiniBand	32 x 400Gb/s	256 x 400Gb/s	72 x 800Gb/s	576 x 800Gb/s
Power Consumption	32 x 1kw=32kw	256 x 1kw=256kw	36 x 2.7kw=97.2kw	288 x 2.7kw=777.6kw
Total Power Consumption			120kw	
Notes	ConnectX-7 NIC		ConnectX-8 NIC	

2.3 NV Blackwell重塑价值链，机柜/铜缆/液冷/HBM占比提升

● GB200 NVL系列的发布，有望带来机柜、HBM、铜缆、液冷等市场的价值量占比提升。

- ✓ **整机柜**：机柜采用MGX架构，由计算托盘与交换托盘组成，提升组装复杂度，带来ODM整机厂的加工价值量提升。
- ✓ **HBM**：H100采用5颗HBM3，Blackwell Ultra预期采用8颗HBM3e，单颗GPU采用HBM数量与单价均实现提升。
- ✓ **铜连接**：GB200 NVL72采用NVLink铜缆链接。
- ✓ **散热**：B200 GPU功耗约1200W，GB200功耗约2700W，或达到风冷上线，有望推动液冷组件价值量提升。

	GB200		H100
 整机柜	\$ 12k 每GPU 服务器ODM毛利情况	2X	\$ 6k
 HBM	\$ 3k 每GPU HBM花费	3X	\$ 1k
 铜连接	\$ 3k 每GPU	10X	\$ 0.3k
 散热	\$ 1.4k 每GPU BOM	3X	\$ 0.4k

第三章 铜连接

DACs市场较快增长，GB200 NVL72需求较大

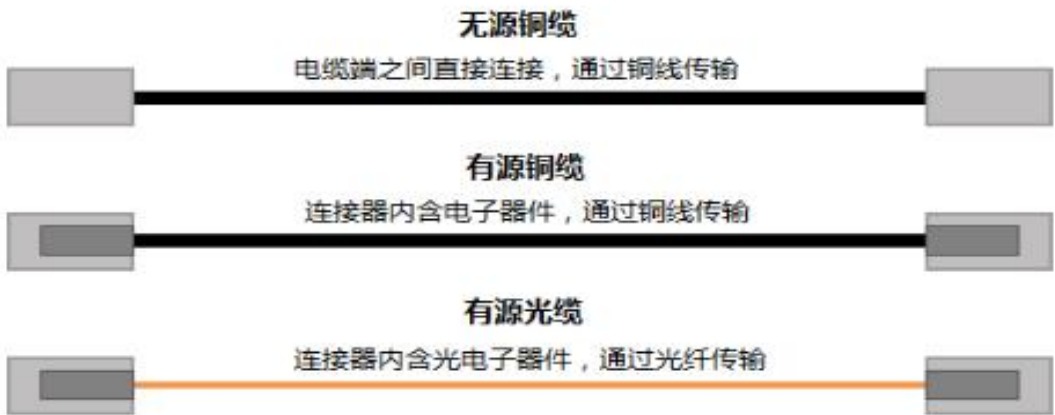
3.1 高速线缆：有源光缆适合远距离传输，直连电缆高速低功耗

- 交换网络中有多种连接解决方案，包括光模块+光纤、有源光缆（AOC）和直连电缆（DAC）。
- 有源光缆AOC：AOC是由集成光电器件组成的，用于数据中心、高性能计算机、大容量存储器等设备间进行高速率、高可靠性互联的传输设备，它通常满足工业标准的电接口，通过内部的电-光-电转换，使用光缆的优越性进行数据的传输。
- 直连电缆DAC：在数据中心，铜缆一般用来连接服务器和存储区域网络，最常用的以直连铜缆（DAC，Direct Attach Copper Cable）为主，而其中又以无源铜缆用得尤盛。由于无源铜缆价格便宜、传输速度快，因而成了实现短距离传输的最佳解决方案。

图：AOC有源光缆

图：DAC直连电缆

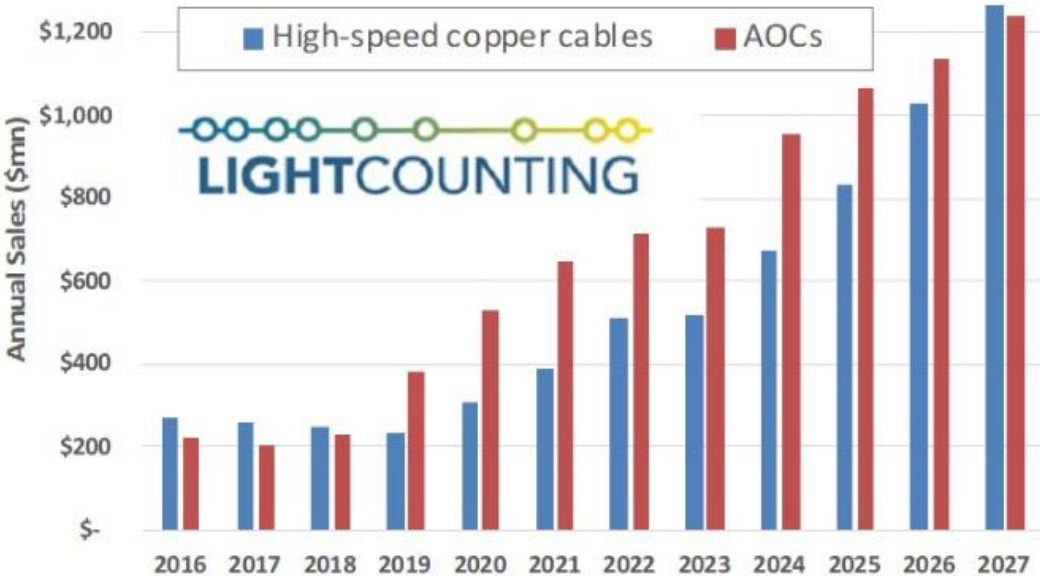
图：高速线缆示意图



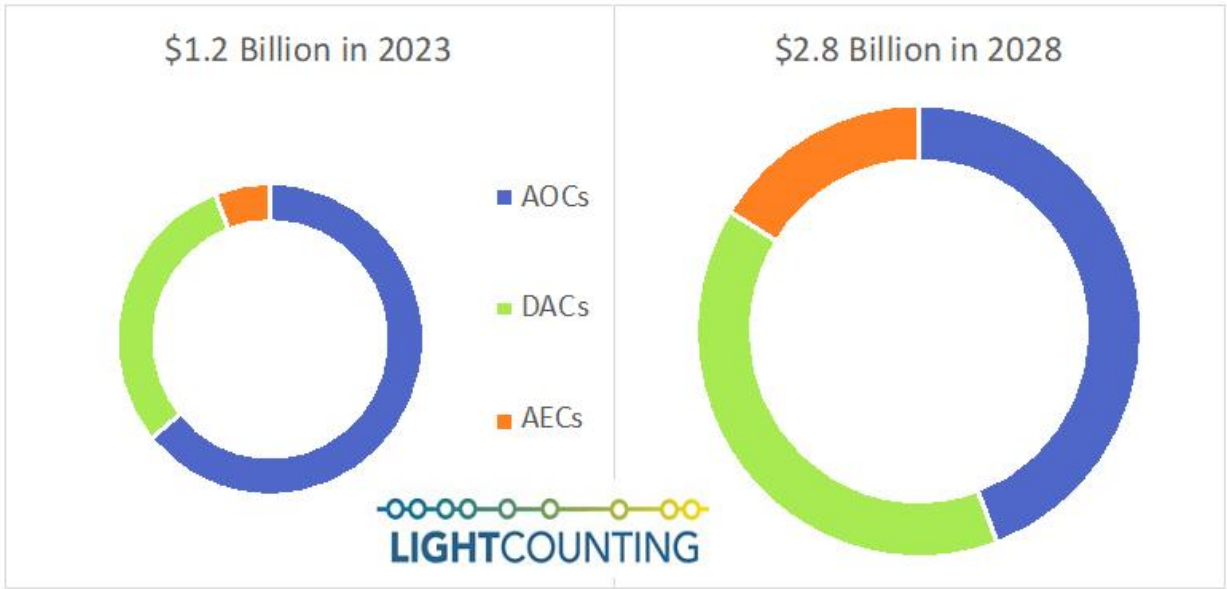
3.1 高速线缆：规模预期28年达28亿美元，DACs保持较快增长

- 据LightCounting，受益于服务器连接以及分解式交换机和路由器中的互连需求提升，预期高速线缆的销售额预计在2023-2028年将增加一倍以上，到2028年将达到28亿美元。有源电缆(AECs)将逐步抢占有源光缆(AOCs)和无源铜线(DACs)的市场份额。AECs的传输距离更长，而且比DACs轻得多。
- DACs将保持较快增长，NVIDIA需求较高。新的100G SerDes使100G / lane设计的无源铜缆的覆盖范围比预期的要长，包括400G、800G和1.6T DACs。据LightCounting，预计2024年底开始发货的200G SerDes表现出色，这将使DACs扩展到每通道200G的设计，由于DACs不需要电，因此它是努力提高电力效率的数据中心连接的默认解决方案。最小化功耗对AI集群来说是最关键的。Nvidia的策略是尽可能多地部署DACs，只有在必要的情况下才使用光模块。

图：AOCs与铜缆的销售量预期



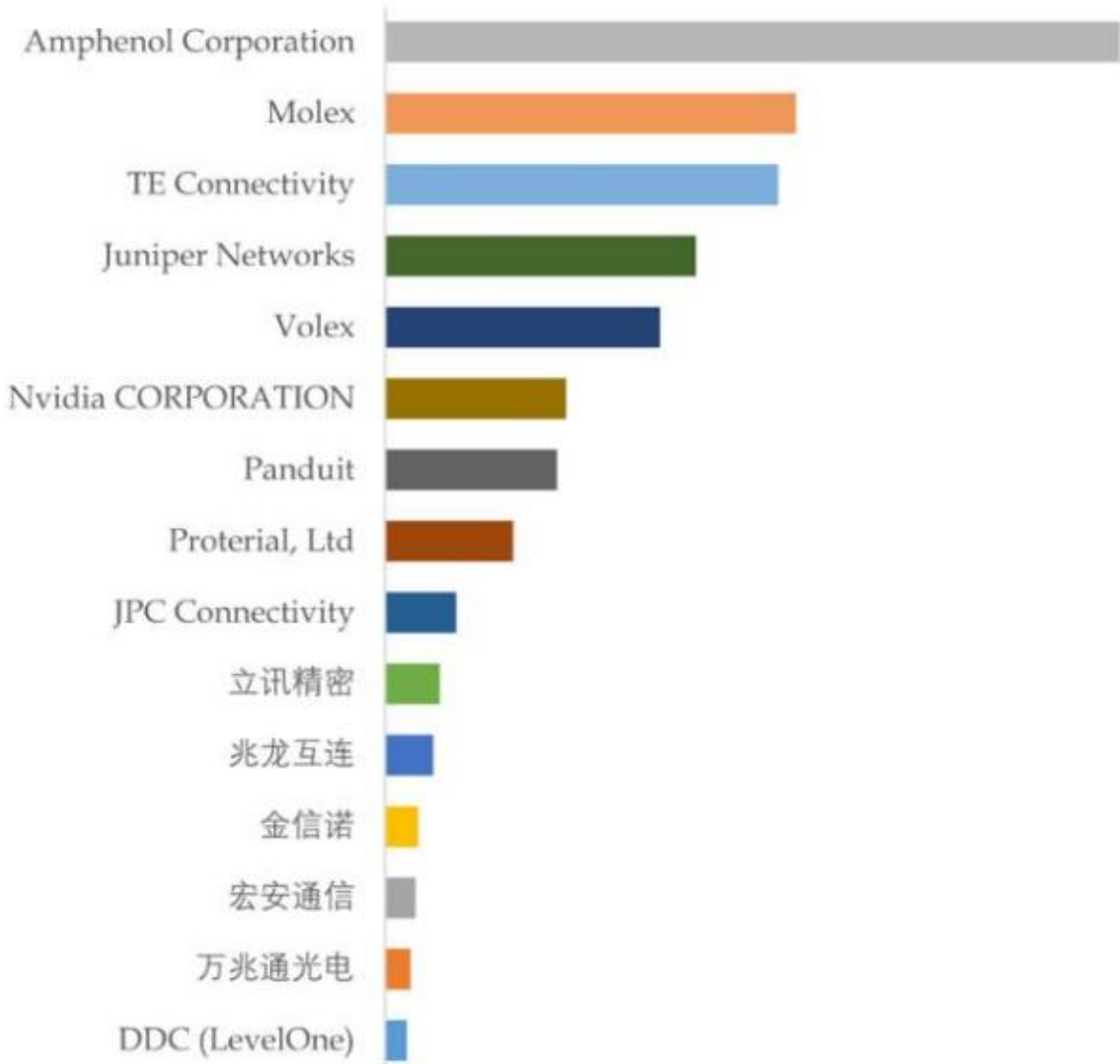
图：AOCs与铜缆的销售量预期



3.1 高速线缆：安费诺为DAC市占率第一，集中度相对较高

- 根据 QYResearch 、北京瑞云智信等数据，在全球范围内，高速直连铜（DAC）电缆的主要生产商有安费诺（Amphenol Corporation）、莫仕（Molex）、泰科（TE Connectivity）、瞻博网络（Juniper Networks）、豪利士（Volex）、泛达公司（Panduit）等。
- 2022 年，全球前十强厂商大约占据了 69.0% 的市场份额，其中安费诺的市场份额位居全球第一。

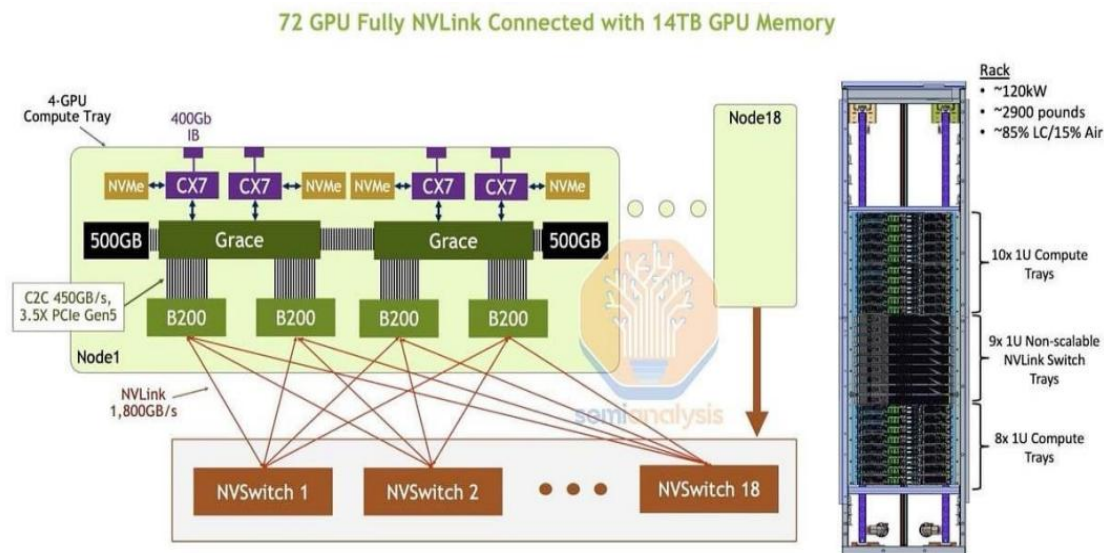
图：2022年，全球高速直连铜缆(DAC)厂商前15强及份额



3.2 英伟达NVL72：采用NVLink全连接，铜缆数量达5000+根

- 据GTC大会，英伟达GB200NVLink Switch和Spine由72个Blackwell GPU采用NVLink全互连，有5000根NVLink铜缆（合计长度超2英里）。
- GB200 NVL72 Switch Tray：**内部有一对 Nvidia 的 NVLink 7.2T ASIC，总共提供 144 个 100 GBps 链路。每个机架有 9 个 NVLink 交换机，相当于为机架中的 72 个 GPU 中的每一个提供 1.8 TBps（18 个链路）的双向带宽。

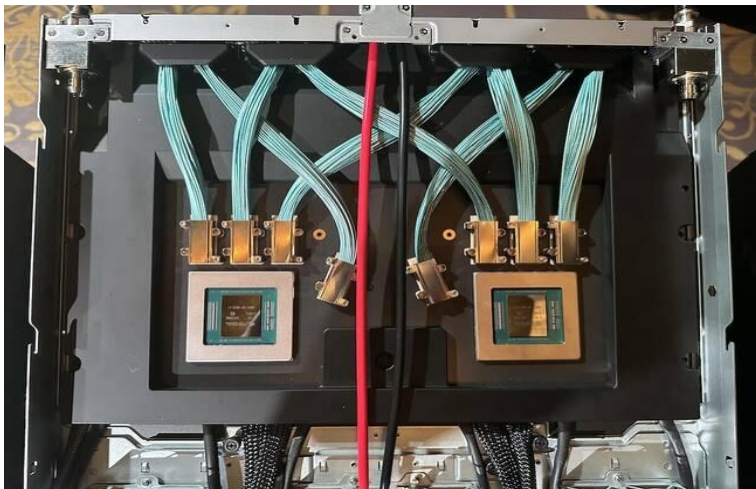
图：英伟达GB200 NVL72 互联结构



图：英伟达GB200 NVL72铜缆组成



图：英伟达GB200 NVL72 Switch Tray



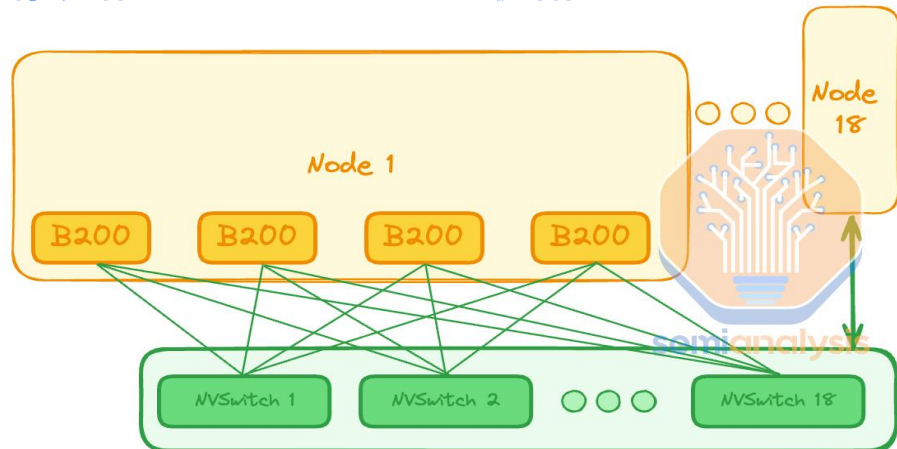
图：英伟达GB200 NVL72 Compute Tray



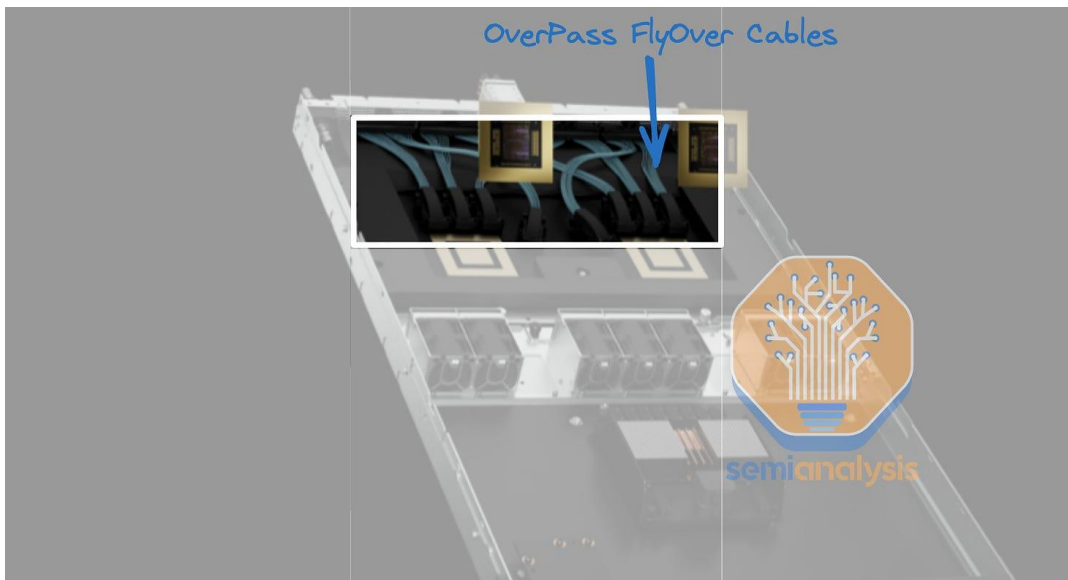
3.2 英伟达NVL72：overpass高速跳线与Ultrapass Paladin背板

- Nvidia 选择使用 Amphenol 的 Ultrapass Paladin 背板产品作为其 NVLink 背板互连的主要初始来源。每个 Blackwell GPU 都连接到 Amphenol Paladin HD 224G/s 连接器，每个连接器有 72 个差分对。然后，该连接器连接到背板 Paladin 连接器。从 NVSwitch ASIC 芯片，需要 OverPass 跨接电缆。

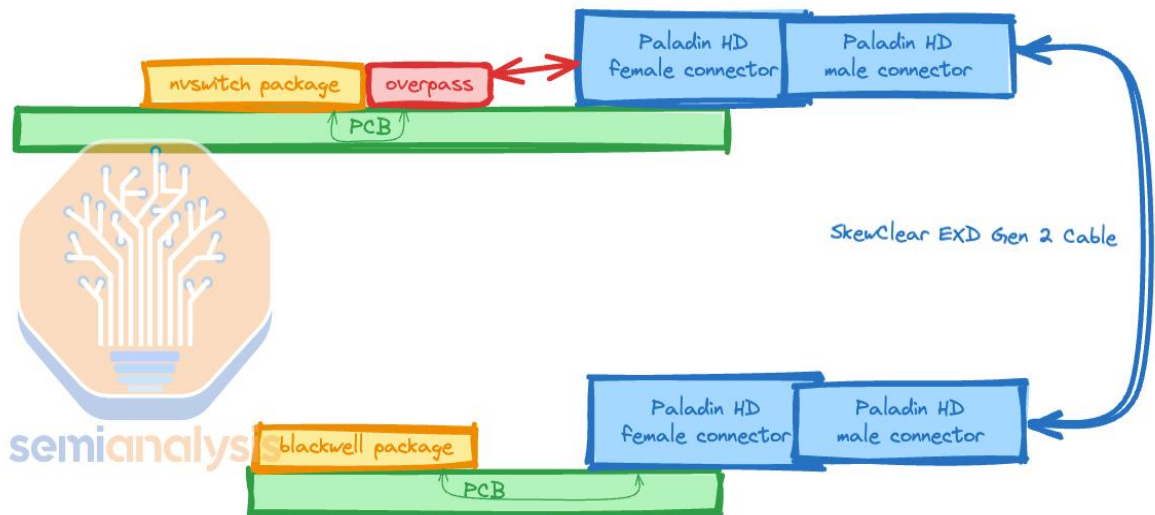
图：英伟达GB200 NVL72 互联结构



图：英伟达GB200 NVL72 Switch Tray



图：英伟达GB200 NVL72 连接系统



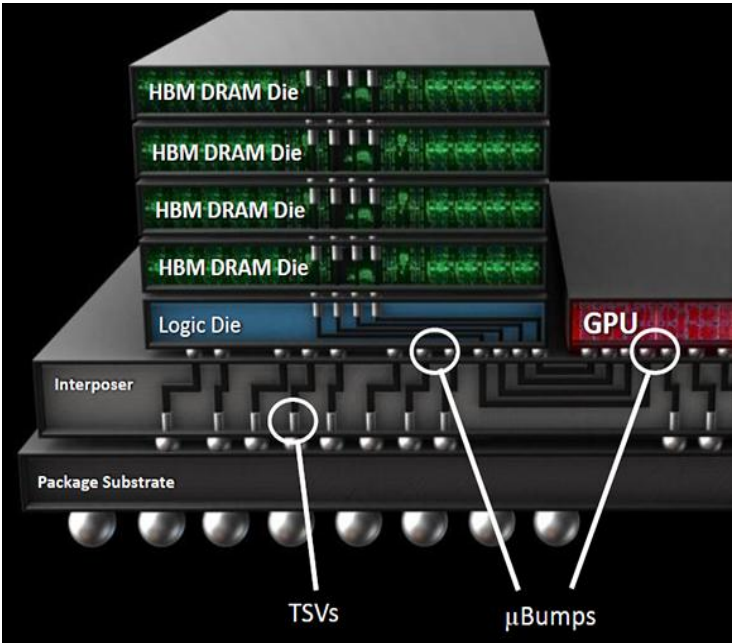
第四章 HBM

HBM3E将于下半年出货，英伟达为主要买家

4.1 HBM市场：HBM3E将于下半年出货，HBM4或于2026年上市

- **HBM (High Bandwidth Memory)**：是一款新型的CPU/GPU 内存芯片（即 “RAM” ），是将很多个DDR芯片堆叠在一起后和GPU封装在一起，实现大容量，高位宽的DDR组合阵列。
- **HBM 目前已经量产的共有 HBM、HBM2、HBM2E、HBM3、HBM3E 五个子代标准。** 从迭代周期看，SK 海力士表示HBM3E内存已经量产，HBM4 产品有望于 2026 年正式推出并量产，同样，美光 and 三星也计划在2026年量产HBM4。

图：HBM DRAM 3D图形



中间的die是GPU/CPU，左右2边4个小die是DDR颗粒的堆叠，Die 之间用TVS方式连接，现在一般只有2/4/8三种层数。

图：海力士的HBM 发展历程

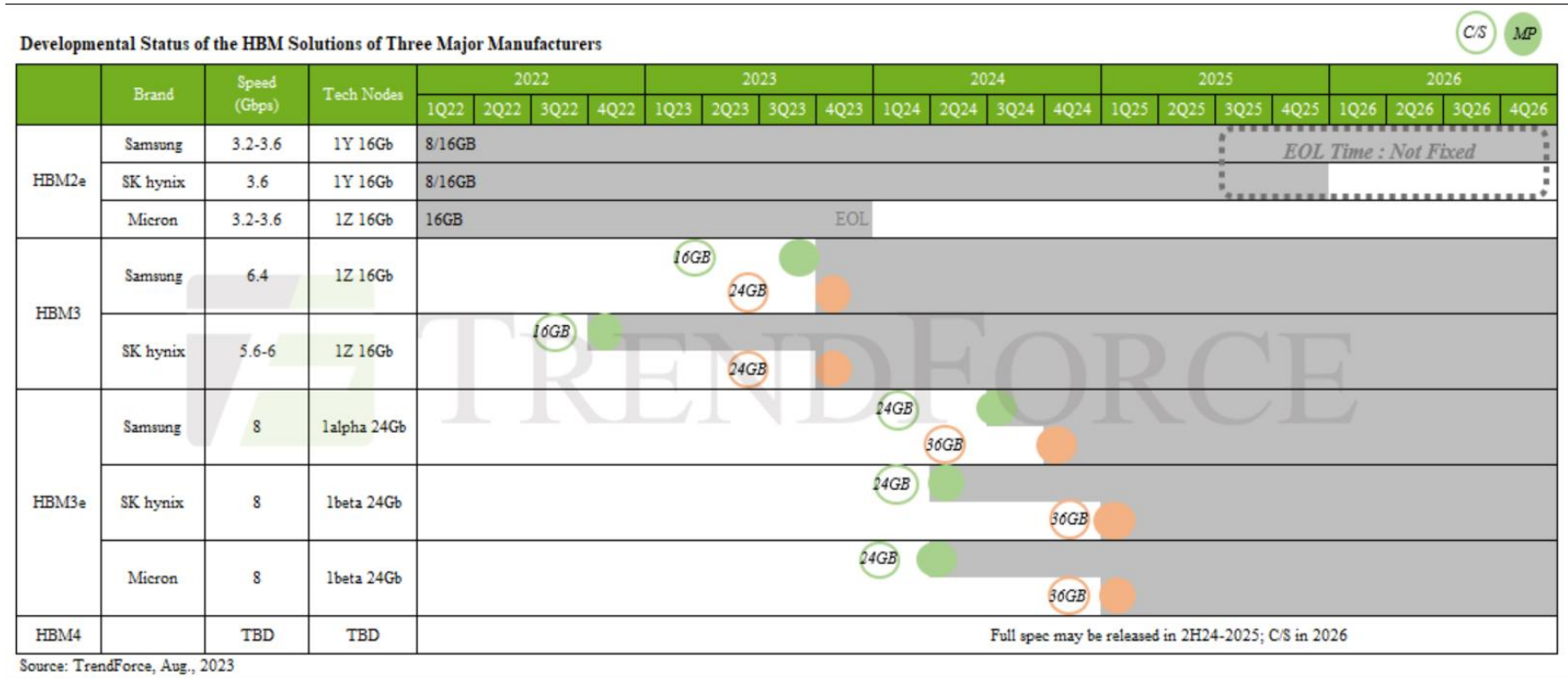


4.1 HBM市场：HBM3E将于下半年出货，HBM4或于2026年上市

● **HBM3e将是今年市场主流，集中在今年下半年出货。**集邦咨询指出，HBM3e将在今年成为市场主流，出货量集中在下半年。目前，SK海力士仍然是主要供应商，与美光一起，都使用1beta nm制程，并且都已开始向英伟达供货。三星使用1alpha nm制程，预计将在第二季度完成认证，于年中开始交付。

● **HBM4有望2026年上市。**据TrendForce集邦，HBM4预计规划于2026年推出。随着客户对运算效能要求的提升，在堆栈的层数上，HBM4除了现有的12hi外，也将再往16hi发展。HBM4 12hi产品将于2026年推出；而16hi产品则预计于2027年问世。此外，受到规格更往高速发展带动，将首次看到HBM最底层的Logic die采用12nm制程wafer，该部分将由晶圆代工厂提供，使得单颗HBM产品需要结合晶圆代工厂与存储器厂的合作。

表：HBM发展时间轴



4.1 HBM：2024年底TSV产能约250K/m，三星与海力士扩产积极



- 据TrendForce预估，截至2024年底，整体DRAM产业规划生产HBM TSV的产能约为250K/m，占总DRAM产能（约1,800K/m）约14%，供给位元年成长约260%。此外，2023年HBM产值占比之于DRAM整体产业约8.4%，至2024年底将扩大至20.1%。
- 以HBM产能来看，三星、SK海力士（SK hynix）至今年底的HBM产能规划最积极，三星HBM总产能至年底将达约130K（含TSV）；SK海力士约120K，但产能会依据验证进度与客户订单持续而有变化。另以现阶段主流产品HBM3产品市占率来看，目前SK海力士于HBM3市场比重逾9成，而三星将随着后续数个季度AMD MI300逐季放量持续紧追。

表：2022~2024年HBM占DRAM产业产值比重预估(百万美元)

	2022	2023	2024E
HBM营收占比	2.6%	8.4%	20.1%
DRAM产业营收	80087	51863	84150

表：2023~2024年各供货商HBM/TSV产能预估

HBM供货商	三星	SK海力士	美光
2023年底HBM TSV产能	45000片/月	45000片/月	3000片/月
2024年底HBM TSV产能	130000片/月	120000-125000片/月	20000片/月

4.2 英伟达目前是HBM最大买家，海力士与三星完成HBM3E验证

- 英伟达目前是HBM市场的最大买家。英伟达H100搭载的HBM容量为80GB，而2024年放量的H200则跃升到144GB。据TrendForce，预期2025年，受到Blackwell平台全面搭载HBM3e、产品层数增加，及单芯片HBM容量的上升，NVIDIA在HBM市场的采购比重将突破70%，对HBM3e市场整体的消耗量或推升至85%以上。Blackwell Ultra或将采用8颗HBM3e 12hi，GB200也有升级可能，再加上B200A的规划，预估2025年12hi产品在HBM3e当中的比重将提升至40%。
- 海力士与三星完成HBM3E验证。据TrendForce，SK hynix（SK海力士）、Micron（美光科技）在2024年第二季开始量产HBM3e。随着进入到HBM3e 12hi阶段，技术难度也提升，验证进度显得更加重要，完成验证的先后顺序可能影响订单的分配比重。据集邦咨询，三星、SK海力士与美光已分别于2024年上半年和第三季提交首批HBM3e 12hi样品，目前处于持续验证阶段。其中SK海力士与美光进度较快，有望于今年底完成验证。

表：HBM产品进度情况

项目	详细内容
海力士	据财联社9/4讯，SK海力士正在将利川厂的M10F产线转向生产HBM，计划明年四季度实现量产。投产后，SK海力士的HBM产能可达每月15万片。设备供应商已收到相关设备的采购订单。 据财联社8/20讯，SK海力士副总裁Ryu Seong-su表示，海力士HBM正受到全球企业的高度关注，苹果、微软、谷歌、亚马逊、英伟达、Meta和特斯拉美股科技股七巨头都向SK海力士提出了定制HBM解决方案的要求。
三星	据财联社9/6讯，三星电子正与台积电合作开发下一代高带宽存储器HBM4人工智能（AI）芯片。 据财联社9/4讯，TrendForce集邦咨询表示，三星电子的HBM3E内存产品已完成验证，并开始正式出货HBM3E 8Hi，主要用于H200，同时Blackwell系列的验证工作也在稳步推进。
美光	据财联社9/9讯，美光表示，目前正在将可供制造的（production-capable）12层第五代HBM（HBM3E）送交给AI产业链重要合作伙伴，以进行验证程序。

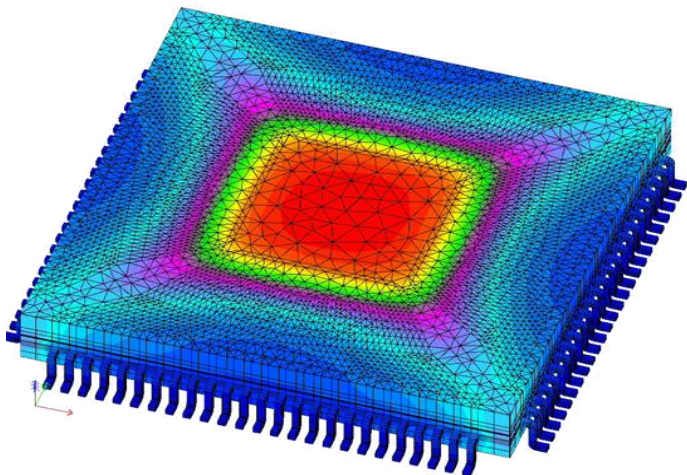
第五章 液冷

冷板式液冷较为成熟，GB200 NVL72采用液冷方案

5.1 芯片散热起源：电子设备发热的本质是工作能量转成热能

- 电子设备发热的本质原因就是工作能量转化为热能的过程。芯片作为电子设备的核心部件，其基本工作原理是将电信号转化为各种功能信号，实现数据处理、存储和传输等功能。而芯片在完成这些功能的过程中，会产生大量热量，这是因为电子信号的传输会伴随电阻、电容、电感等能量损耗，这些损耗会被转化为热能。
- 温度过高会影响电子设备工作性能，甚至导致电子设备损坏。据《电子芯片散热技术的研究现状及发展前景》，如对于稳定持续工作的电子芯片，最高温度不能超过85 °C，温度过高会导致芯片损坏。
- 散热技术需要持续升级，来控制电子设备的运行温度。芯片性能持续发展，这提升了芯片功耗，也对散热技术提出了更高的要求。此外，AI大模型的训练与推理需求，要求AI芯片的单卡算力提升，有望进一步打开先进散热技术的成长空间。

图：芯片的温度云图变化



图：CPU\GPU\Switch IC的TDP逐渐提升

CPU	型号	发布时间	TDP	GPU	型号	发布时间	TDP
Intel	Ice Lake	2021Q2	105W-270W	NVIDIA	A100	2020	250W-400W
	Sapphire Rapids	2023	115W-350W		H100 SXM H100 PCIe	2023	300W-350W up to 700W
	Emerald Rapids	2023	150W-385W	AMD	MI100	2020	up to 300W
	Granite Rapids	2024	500W		MI200	2021	300W
AMD	Rome	2019	120W-280W		MI300	2023	600W
	Milan	2021	120W-280W				
	Genoa	2022	200W-360W				

5.1 芯片散热革新：相比热管/VC/3DVC，冷板式散热范围广

- 热管与VC的散热能力较低，3D VC风冷散热上限扩大至1000W，冷板式具备1000W+的广阔散热范围。



5.1 芯片散热革新：浸没式散热效果好，冷板式更为成熟

- 根据ODCC《冷板液冷服务器设计白皮书》，综合考量初始投资成本、可维护性、PUE 效果以及产业成熟度等因素，冷板式和单相浸没式相较其他液冷技术更有优势，是当前业界的主流解决方案。

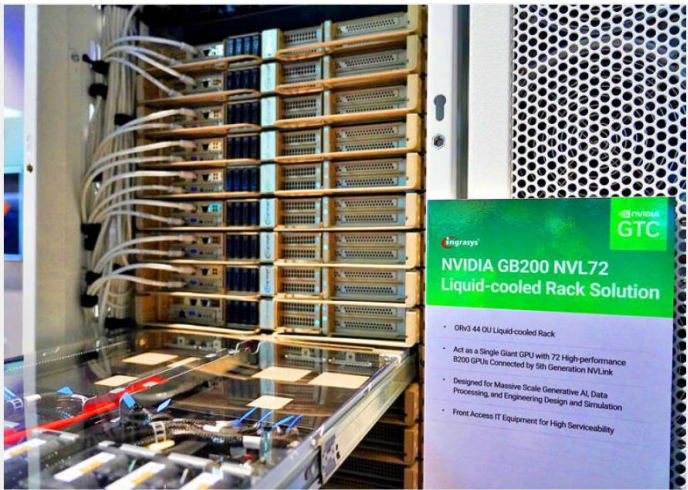
表：服务器散热技术对比情况

液冷方案	非接触式液冷	接触式液冷	
	冷板式	浸没式液冷	
		相变浸没式	单相浸没式
投资成本	初始投资中等 运维成本低	初始投资及运维成本高	初始投资及运维成本高
PUE	1.1-1.2	<1.05	<1.09
可维护性	较简单	复杂	
供应商	华为、浪潮曙光、联想超聚变等主流供应商	曙光	阿里巴巴、H3C、绿色云图、云酷智能、曙光数创
应用案例	多	超算领域较多	较多
分析	初始投资中等，运维成本较低，PUE 收益中等，部署方式与风冷相同，从传统模式过渡较平滑	初始投资最高，PUE 收益最高，使用专用机柜，服务器结构需改造为刀片式	初始投资较高，PUE 收益较高，部分部件不兼容，服务器结构需改造

5.2 GB200 NVL72：采用冷板液冷散热，多家厂商合作打造系统

- 在3月18日英伟达GTC2024上，多家厂商展出基于英伟达GB200 NVL72液冷服务器设备，选取冷板液冷进行散热。
- 鸿海：展出Ingrasys鸿佰科技开发的GB200 NVL72液冷机架系统，采用冷板式结构，机架具备风液混合、液冷等结构。

图：鸿海采用液体冷却的 GB200 计算托盘



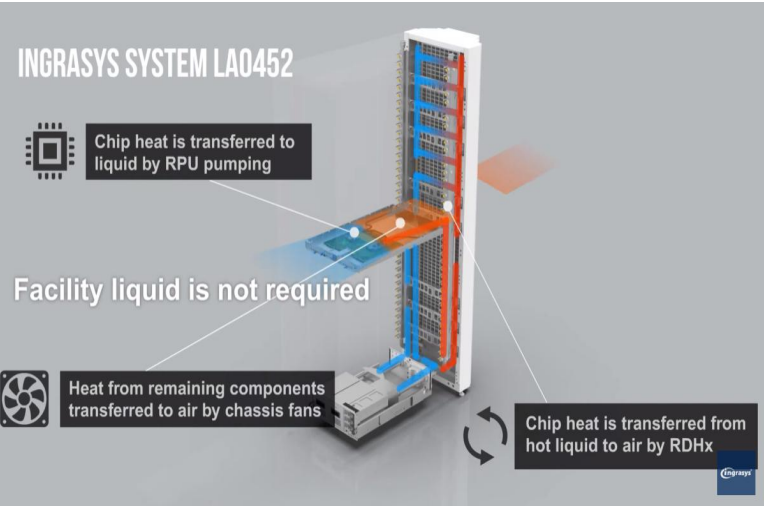
图：英伟达NV72液冷机柜设计



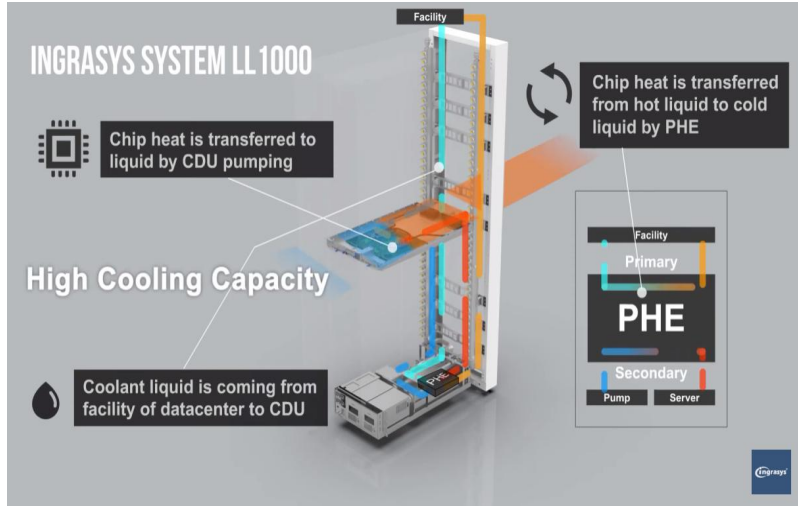
图：鸿海的3种液冷机架解决方案



图：鸿海LA0452系统为液冷-风冷混合散热



图：鸿海LL1000系统为液冷与CDU等设备混合散热



第六章

投资建议与风险提示

- 投资建议：大模型训推带动 AI算力需求增长，GB200等新一代算力架构将推出，算力产业链中的AI芯片、服务器整机、铜连接、HBM、液冷、光模块、IDC等环节有望持续受益。维持对计算机行业“推荐”评级。

- 相关公司

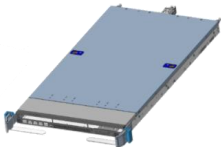
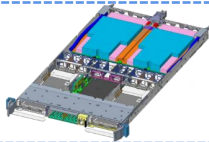
- ✓ 1) AI芯片：海光信息、寒武纪、龙芯中科、景嘉微、英伟达、AMD、Intel
- ✓ 2) 服务器整机：工业富联、浪潮信息、中科曙光、华勤技术、中国长城、高新发展、神州数码、烽火通信、拓维信息、纬创、广达、英业达、纬颖、超微电脑。
- ✓ 3) 服务器组件：①散热：飞荣达、曙光数创、英维克、同方股份、申菱环境、高澜股份、奇鋆科技、双鸿、VERTIV；②主板：沪电股份、深南电路、胜宏科技、技嘉、华擎；③HBM：SK海力士、三星、美光、赛腾股份、联瑞新材；④铜连接：安费诺、沃尔核材、华丰科技。
- ✓ 4) 光模块：天孚通信、中际旭创、新易盛、光迅科技、华工科技。
- ✓ 5) 数据中心：奥飞数据、光环新网、宝信软件、数据港、电科数字。

6.1 投资建议与相关公司



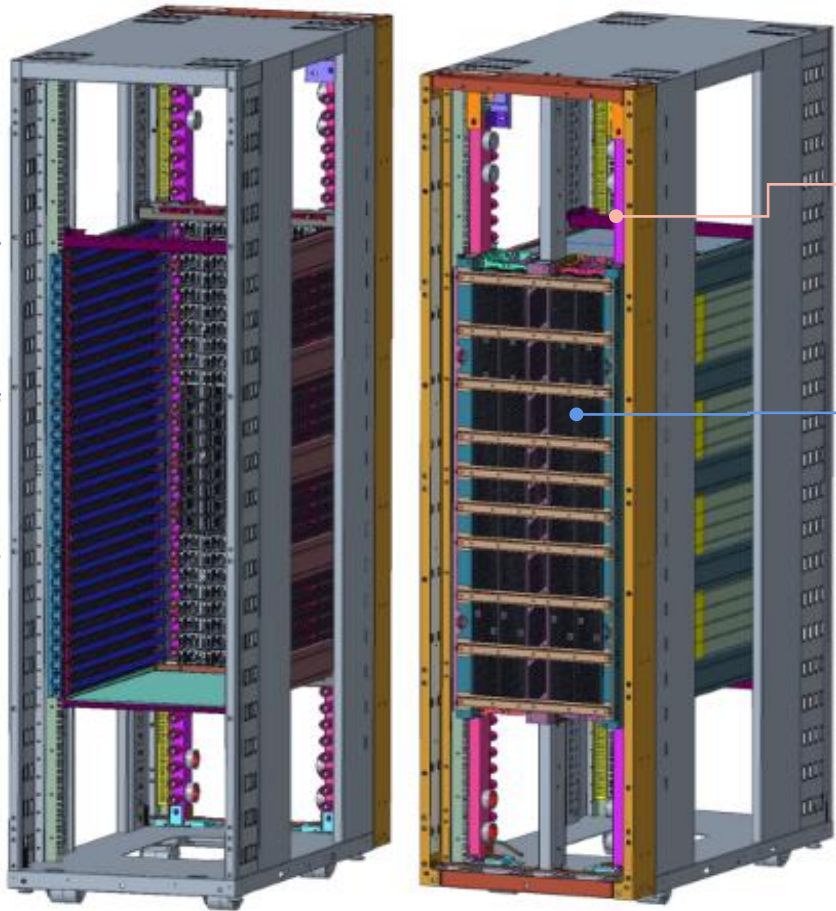
英伟达GB200 NVL72实物图

- **电源**
 - 内地：中国长城、欧陆通、杰华特、泰嘉股份
 - 中国台湾：台达
- **☆GPU**
 - 内地：华为海思、海光信息、寒武纪、龙芯中科。
 - 海外：英伟达、AMD、Intel
- **☆CPU**
 - 内地：华为海思、海光信息、飞腾、龙芯中科
 - 海外：英伟达、AMD、Intel
- **☆主板**
 - 内地：沪电股份、胜宏科技、深南电路、生益科技
 - 中国台湾：泰安电脑（神达）、技嘉、华擎、华硕、英业达、纬创、研华
 - 美国：Intel、Supermicro
- **Compute Tray * 10**
 - 内地：工业富联
 - 中国台湾：纬创
- **Switch Tray * 9**
 - 内地：工业富联
 - 中国台湾：纬创
- **Compute Tray * 8**
- **NV Switch Chip\网卡等**
 - 海外：英伟达、Mellanox
- **电源**



- **存储 NAND**
 - 公司：长江存储、兆易创新、佰维存储、朗科科技
 - 海外：三星、西部数据、铠侠、SK海力士、美光

- **光模块**
 - 内地：中际旭创、新易盛、光迅科技、天孚通信



正面

背面

- **☆分歧管（液冷：支持130KW的制冷能力）**
 - 内地：曙光数创、飞荣达、中航光电、英维克、同飞股份、申菱环境、高澜股份、依米康
 - 中国台湾：奇铤、双鸿、健策、建准、高力
 - 海外：VERTIV

- **☆线缆（铜连接：5000+根NVLink线缆）**
 - 内地：沃尔核材、华丰科技、兆龙互连、鼎通科技、立讯精密、神宇股份
 - 海外：安费诺

- **☆组装测试**
 - 整机厂商ODM：1) 内地：浪潮信息、工业富联、中科曙光、紫光股份、中国长城、华勤技术、神州数码、拓维信息、烽火通信、软通动力；2) 中国台湾：鸿海、广达、纬创、英业达、纬颖；3) 美股：戴尔、超微电脑
 - BIOS/BMC：1) 内地：卓易信息；2) 中国台湾：新唐、系微、信骅

注：蓝字为零部件。上图展示以英伟达GB200 NVL72整机柜为例，仅列示了部分参与公司

- 1) 宏观经济影响下游需求：宏观经济环境下行，将影响客户对信息化基础设施的采购需求；
- 2) 大模型产业发展不及预期：行业的核心驱动力是人工智能大语言模型的训练和推理对算力的需求，如果大模型行业发展不及预期将会影响AI算力的相关需求；
- 3) 市场竞争加剧：IT 产品和服务行业是成熟且完全竞争的行业，新进入者可能加剧整个行业的竞争态势；
- 4) 中美博弈加剧：国际形势持续不明朗，美国不断通过“实体清单”等方式对中国企业实施打压，若中美紧张形势进一步升级，将可能导致中国半导体供应链供应受到影响；
- 5) 相关公司业绩不及预期：市场环境变化、公司治理情况变化、其他非主营业务经营不及预期等原因或将导致相关公司的整体业绩不及预期。
- 6) 各公司并不具备完全可比性，对标的相关资料和数据仅供参考。

计算机小组介绍

刘熹，计算机行业首席分析师，上海交通大学硕士，多年计算机行业研究经验，致力于做前瞻性深度研究，挖掘投资机会。新浪金麒麟新锐分析师、金牌分析师团队核心成员。

分析师承诺

刘熹，本报告中的分析师均具有中国证券业协会授予的证券投资咨询执业资格并注册为证券分析师，以勤勉的职业态度，独立，客观的出具本报告。本报告清晰准确的反映了分析师本人的研究观点。分析师本人不曾因，不因，也将不会因本报告中的具体推荐意见或观点而直接或间接收取到任何形式的补偿。

国海证券投资评级标准

行业投资评级

推荐：行业基本面向好，行业指数领先沪深300指数；
中性：行业基本面稳定，行业指数跟随沪深300指数；
回避：行业基本面向淡，行业指数落后沪深300指数。

股票投资评级

买入：相对沪深300 指数涨幅20%以上；
增持：相对沪深300 指数涨幅介于10%~20%之间；
中性：相对沪深300 指数涨幅介于-10%~10%之间；
卖出：相对沪深300 指数跌幅10%以上。

免责声明

本报告的风险等级定级为R3，仅供符合国海证券股份有限公司（简称“本公司”）投资者适当性管理要求的客户（简称“客户”）使用。本公司不会因接收人收到本报告而视其为客户。客户及/或投资者应当认识到有关本报告的短信提示、电话推荐等只是研究观点的简要沟通，需以本公司的完整报告为准，本公司接受客户的后续问询。

本公司具有中国证监会许可的证券投资咨询业务资格。本报告中的信息均来源于公开资料及合法获得的相关内部外部报告资料，本公司对这些信息的准确性及完整性不作任何保证，也不保证其中的信息已做最新变更，也不保证相关的建议不会发生任何变更。本报告所载的资料、意见及推测仅反映本公司于发布本报告当日的判断，本报告所指的证券或投资标的的价格、价值及投资收入可能会波动。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。报告中的内容和意见仅供参考，在任何情况下，本报告中所表达的意见并不构成对所述证券买卖的出价和征价。本公司及其本公司员工对使用本报告及其内容所引发的任何直接或间接损失概不负责。本公司或关联机构可能会持有报告中所提到的公司所发行的证券头寸并进行交易，还可能为这些公司提供或争取提供投资银行、财务顾问或者金融产品等服务。本公司在知晓范围内依法合规地履行披露义务。

风险提示

市场有风险，投资需谨慎。投资者不应将本报告为作出投资决策的唯一参考因素，亦不应认为本报告可以取代自己的判断。在决定投资前，如有需要，投资者务必向本公司或其他专业人士咨询并谨慎决策。在任何情况下，本报告中的信息或所表述的意见均不构成对任何人的投资建议。投资者务必注意，其据此做出的任何投资决策与本公司、本公司员工或者关联机构无关。

若本公司以外的其他机构（以下简称“该机构”）发送本报告，则由该机构独自为此发送行为负责。通过此途径获得本报告的投资者应自行联系该机构以要求获悉更详细信息。本报告不构成本公司向该机构之客户提供的投资建议。

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。本公司、本公司员工或者关联机构亦不为该机构之客户因使用本报告或报告所载内容引起的任何损失承担任何责任。

郑重声明

本报告版权归国海证券所有。未经本公司的明确书面特别授权或协议约定，除法律规定的情况外，任何人不得对本报告的任何内容进行发布、复制、编辑、改编、转载、播放、展示或以其他方式非法使用本报告的部分或者全部内容，否则均构成对本公司版权的侵害，本公司有权依法追究其法律责任。

国海证券 · 研究所 · 计算机研究团队

心怀家国，洞悉四海



国海研究上海

上海市黄浦区绿地外滩中心C1栋
国海证券大厦

邮编：200023

电话：021-61981300

国海研究深圳

深圳市福田区竹子林四路光大银行大厦28F

邮编：518041

电话：0755-83706353

国海研究北京

北京市海淀区西直门外大街168号腾达大厦25F

邮编：100044

电话：010-88576597