

推理算力重要性提升，光模块发展再加速

2025 年 03 月 23 日

核心观点

- **事件：**近日英伟达 GPU 技术大会 2025（GTC 2025）结束。
- **推理算力预计将持续增长，Scaling Law 仍将持续：**会上，英伟达 CEO 表示，通往 AGI 以及具身智能机器人均需要算力，构建 Omniverse 与世界模型更需要源源不断的算力，至于最终人类构建一个虚拟的“平行宇宙”，英伟达认为算力需求将是过去的 100 倍。为了支撑该观点，黄仁勋列出云厂商采购数据，即 2024 年美国前四云厂总计采购 130 万颗 Hopper 架构芯片，到了 2025 年，这一数据增至 360 万颗 Blackwell GPU，成倍数增长。根据 Semianalysis 数据，如今模型需要处理超过 100 万亿个 token，推理模型的 token 数量是之前的 20 倍，计算量是之前的 150 倍。同时，英伟达 CEO 黄仁勋进一步指出，Blackwell 的性能比 Hopper 提高了 68 倍，导致成本下降了 87%。Rubin 计划推动更多的性能提升——是 Hopper 的 900 倍，成本降低 99.97%，算法带来的效率提升与硬件迭代带来的性能提升共振。

- **硬件端对推理算力的侧重程度更强，软件端发布 Dynamo 及 Llama Nemotron，英伟达 AI Agent 发展动能充分：**我们认为此次 GTC 2025，相较于此前 GTC 大会，硬件性能提升符合预期，英伟达对软件的布局、推理算力以及 Agent 的愿景更强，硬件发展带动软件及大模型持续高景气。

硬件方面：英伟达发布基于 Blackwell 架构的升级版 Blackwell Ultra，并着重强调了其在推理端的重大进展，可以为数据中心提供 50 倍增收的机会；英伟达也明确了 26-27 年乃至更远期的发展规划，硬件性能进一步加速已成定局；CPO 交换机方面，英伟达展示了基于 1.6T 硅光引擎的 CPO（共封装光学）交换机系列，硬件端性能未来仍将持续高质量全方位提升。

软件方面：英伟达发布了专注于简化推理部署和扩展的开放 AI 引擎栈 Nvidia Dynamo，可能成功开创推理软硬件效率的新范式；同时发布 Nvidia Llama Nemotron，望以 Llama Nemotron 推理模型作为基础，通过 NVIDIA AIQ 架构作为 Agent 模板进行相关方向探索，形成生态闭环。

- **投资建议：**当前算力板块在推理应用的不断发展中并非呈现出算力需求下降的态势，而是算力需求得到进一步的刺激。当前市场普遍认为推理应用的不断增长将挤压 2026 年甚至更长时间段内算力板块的成长空间，我们认为该种判断可能同算力行业的实际情况有所偏离。在算力需求总体增长的情况下，我们认为当前算力相关板块仍然具备较大投资价值，建议优选空间较大且动能充足的细分子板块：

建议关注：运营商中国移动、中国联通、中国电信等；光模块中际旭创、新易盛、天孚通信、光迅科技、华工科技等；光芯片源杰科技、仕佳光子等；AIDC 相关光环新网、数据港等。

- **风险提示：**AI 发展不及预期的风险，全球地缘政治不确定性的风险，算力发展不及预期的风险等。

通信行业

推荐 维持评级

分析师

赵良毕

☎：010-8092-7619

✉：zhaoliangbi_yj@chinastock.com.cn

分析师登记编码：S0130522030003

相对沪深 300 表现图

2025-03-23



资料来源：Wind，中国银河证券研究院

相关研究

1. 【银河通信】深度报告：DeepSeek 冲击波，通信算力降本增效，Deepseek 变与不变
2. 【银河通信】2025 年度策略报告：高成长高景气，科技变革创长牛
3. 【银河通信】科技行业专题报告：大拐点，大机遇，创新变革，拥抱科技新动能
4. 【银河通信】运营商行业 2024 年中报专题：运营商利润增速稳健，数智化转型全球领先
5. 【银河通信】出海专题报告：技术+成本优势驱动，市占率持续提升
6. 【银河通信】中期策略报告：AI 为基算力为石，科技变革浩瀚星辰

目录

Catalog

一、 GTC 2025：硬件性能再提升，软件高速发展	3
（一） 英伟达推出 Blackwell Ultra 架构，AI 工厂推理能力进一步加强.....	3
（二） 1.6T 硅光引擎 CPO 交换机推出，光电器件有望 2H25 量产	5
（三） Nvidia Dynamo 及 Llama Nemotron 发布，软件层面持续发力	6
二、 投资建议	8
三、 风险提示	9

一、GTC 2025：硬件性能再提升，软件高速发展

英伟达 GPU 技术大会 (GTC, 2025 年 3 月 17 日-3 月 21 日) 结束, 本次大会分为多个环节, 最引人注目的环节为 2025 年 3 月 19 日英伟达 CEO 黄仁勋进行的主题演讲《AI 工厂》, 详细的介绍了英伟达 Blackwell 系列芯片全面投产、旗舰产品更新规划、个人端 AI 超算、量子计算、硅光网络交换机以及通用机器人等主题。

推理算力预计将持续增长, Scaling Law 仍将持续:会上, 英伟达 CEO 黄仁勋表示, 通往 AGI 需要算力, 具身智能机器人需要算力, 构建 Omniverse 与世界模型更需要源源不断的算力, 至于最终人类构建一个虚拟的“平行宇宙”, 英伟达认为算力需求将是过去的 100 倍。为了支撑该观点, 黄仁勋列出云厂商采购数据, 即 2024 年美国前四云厂总计采购 130 万颗 Hopper 架构芯片, 到了 2025 年, 这一数据增至 360 万颗 Blackwell GPU, 成倍数增长。根据 Semianalysis 数据, 如今模型需要处理超过 100 万亿个 token, 推理模型的 token 数量是之前的 20 倍, 计算量是之前的 150 倍。同时, 英伟达 CEO 黄仁勋进一步指出, Blackwell 的性能比 Hopper 提高了 68 倍, 导致成本下降了 87%。Rubin 计划推动更多的性能提升——是 Hopper 的 900 倍, 成本降低 99.97%, 算法带来的效率提升与硬件迭代带来的性能提升共振。

(一) 英伟达推出 Blackwell Ultra 架构, AI 工厂推理能力进一步加强

Blackwell Ultra 架构发布, HBM 内存容量升级, 显存提升带动计算性能增长。英伟达于 2024 年 3 月发布 Blackwell 架构并推出 GB200 芯片, 2025 年则推出 Blackwell Ultra 架构, 并将于 2025 年下半年上市, Blackwell Ultra 延续此前 Blackwell 基础架构, 包括双芯片设计和高速 NVLink 互联技术由两颗台积电 N4P (5nm) 工艺, Blackwell 架构芯片+Grace CPU 封装而来, 并且搭配了比之前 GB200 芯片更先进的 12 层堆叠的 HBM3e 内存, 显存提升为 288GB, 和上一代一样支持第五代 NVLink, 可实现 1.8TB/s 的片间互联带宽。Blackwell Ultra 增强了训练和测试时间扩展推理能力, 使各地组织能够加速人工智能推理、代理人工智能和物理人工智能等应用, 黄仁勋表示: “人工智能已经取得了巨大的飞跃--推理和代理式人工智能需要更多数量级的计算性能。

“我们为这一时刻设计了 Blackwell Ultra--它是一个单一的多功能平台, 能够轻松高效地进行前训练、后训练和推理 AI 推断。”

表1: NVLink 历代性能参数对比及进化

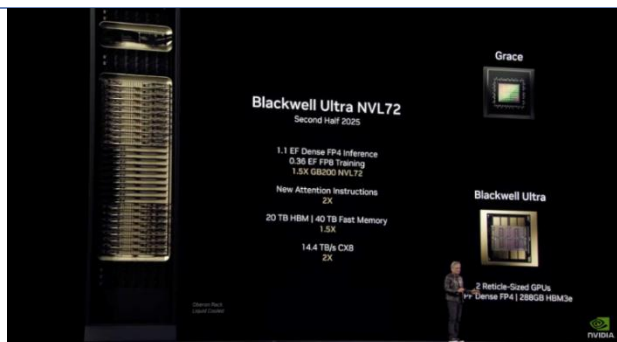
	NVLink V1	NVLink V2	NVLink V3	NVLink V4	NVLink V5
带宽	160GB/s	300GB/s	600GB/s	900GB/s	1.8T/s
应用产品	P100	V100	A100	H200	GB200

资料来源: 英伟达, 芯智讯, 中国银河证券研究院

基于 Blackwell Ultra 推出的 Blackwell Ultra NVL72 机柜主要用于推理任务: Blackwell Ultra NVL72 机柜一共由 18 个计算托盘构成, 每个计算托盘包含 4 颗 Blackwell Ultra GPU+2 颗 Grace CPU, 总计 72 颗 Blackwell Ultra GPU+36 颗 Grace CPU, HBM 容量达到 20TB, 总带宽 576TB/s, 外加 9 个 NVLink 交换机托盘 (18 颗 NVLink 交换机芯片), 节点间 NVLink 带宽 130TB/s。机柜内置 72 张 CX-8 网卡, 提供 14.4TB/s 带宽, Quantum-X800 InfiniBand 和 Spectrum-X 800G 以太网卡则可以降低延迟和抖动, 支持大规模 AI 集群。此外, 机架还整合了 18 张用于增强多租户网络、安全性和数据加速 BlueField-3 DPU。英伟达说这款产品是“为 AI 推理时代”专门定制, 在进行 FP4 精度的推理任务时, 能够达到 1.1 ExaFLOPS (每秒百亿亿次浮点运算); 在进行 FP8 精度的训练任务时, 性能为 1.2 ExaFLOPS。相比前一代产品 GB200 NVL72 的 AI 性能提

升到了 1.5 倍，HBM 容量也提升到了 1.5 倍，支持的 40TB 快速内存容量也是前代的 1.5 倍，网卡总带宽是前代的 2 倍。相比 Hopper 架构同定位的 DGX 机柜产品，可以为数据中心提供 50 倍增收的机会。应用场景包括推理型 AI、Agent 以及物理 AI(用于机器人、智驾训练用的数据仿真合成)。根据英伟达相关信息,6710 亿参数 DeepSeek-R1 的推理,基于 H100 产品可实现每秒 100tokens,而采用 Blackwell Ultra NVL72 方案,可以达到每秒 1000 tokens。换算成时间,同样的推理任务,H100 需要跑 1.5 分钟,而 Blackwell Ultra NVL72 15 秒即可跑完。

图1: Blackwell Ultra NVL72-AI 推理为主



资料来源: Nvidia, 中国银河证券研究院

图2: Blackwell Ultra NVL72 同上代产品对比

	Blackwell Ultra NVL72	GB200 NVL72
Blackwell芯片 Grace芯片	72 36	72 36
FP4	1.1ExaFLOPS	1440PetaFLOPS
快速内存	40TB	30TB
内存	20TB	14TB
内存带宽	576TB	576TB/s
NVLink带宽	130TB/s	130TB/s

资料来源: 芯智讯, 中国银河证券研究院

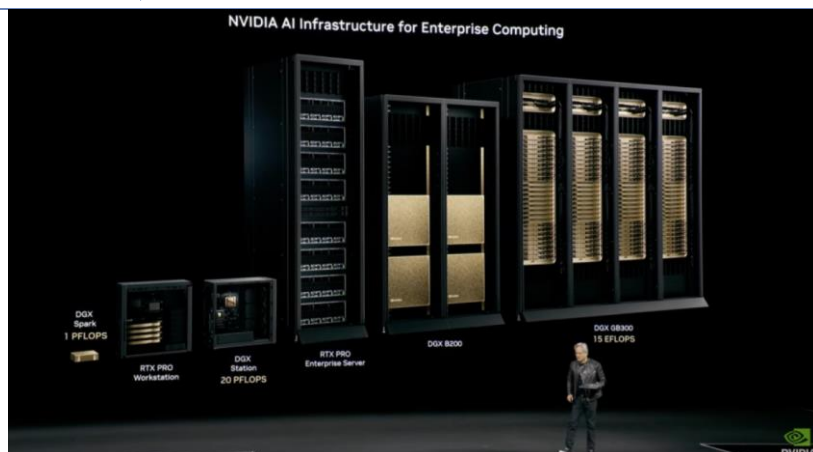
AI 工厂：基于 Blackwell Ultra 定制的 DGX Super POD 包括 GB300（1*Grace CPU+2*Blackwell Ultra GPU） NVL72 机架级解决方案（液冷）和 NVIDIA HGX B300（不含 Grace CPU） NVL16 系统（风冷）两种方案。AI 工厂为代理式 AI、生成式 AI 和物理 AI 工作负载提供专用基础设施，此类工作负载需要海量计算资源，才能为在生产环境中运行的应用提供 AI 预训练、后训练和测试阶段扩展功能。其中：

DGX GB300 系统：采用 NVIDIA Grace Blackwell Ultra 超级芯片（搭载 36 个 NVIDIA Grace™ CPU 和 72 个 NVIDIA Blackwell Ultra GPU），并采用机架级、液冷架构，专为实现高级推理模型的实时代理响应而设计，而 DGX SuperPOD 则搭载总计 288 颗 Grace CPU+576 颗 Blackwell Ultra GPU，提供 300TB 的快速内存，FP4 精度下算力为 11.5ExaFLOPS；

风冷 NVIDIA DGX B300 系统：采用 NVIDIA B300 NVL16 架构，不含 Grace CPU 芯片，具备相对较强的扩展空间，助力全球各地的数据中心满足生成式和代理式 AI 应用的计算需求，主要针对普通的企业级数据中心。

同时，英伟达在本次的 GTC 大会上也进一步强化了 Blackwell 和 RTX 系列的整合，推出了内置 GDDR7 内存的 AI PC 相关 GPU，覆盖笔记本、桌面甚至是数据中心等场景，进一步完善产品矩阵。

图3: 英伟达 RTX 与 DGX 产品矩阵融合，覆盖场景进一步加强



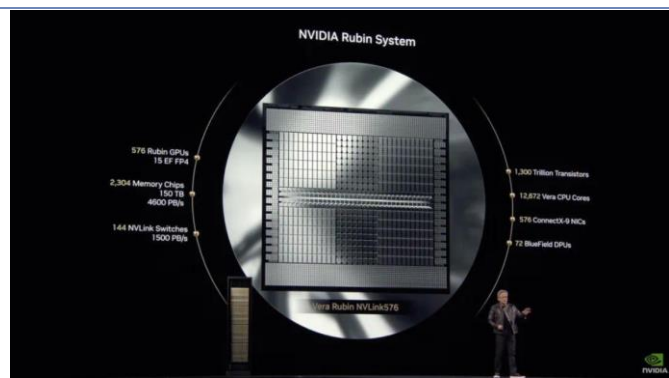
资料来源: 英伟达, 芯智讯, 中国银河证券研究院

本届 GTC 大会中，英伟达 CEO 黄仁勋也对未来几年的 GPU 发展规划进行分享：

2026 年将上市基于 Rubin 架构的下一代 GPU 以及更强的机柜 Vera Rubin NVL144。在硬件配置上，Rubin Ultra 的 Veras 系统延续了 88 个定制 Arm 核心的设计，每个核心支持 176 个线程，并通过 NVLink-C2C 提供 1.8 TB/s 的带宽。而 GPU 方面，Rubin Ultra 集成了 4 个 Reticle-Sized GPU，每颗 GPU 提供 100 petaflops 的 FP4 计算能力，并配备 1TB 的 HBM4e 内存，在性能和内存容量上都达到了新的高度。Vera Rubin NVL144 则集成的了 72 颗 Vera CPU+144 颗 Rubin GPU，采用 288GB 显存的 HBM4 芯片，显存带宽 13TB/s，搭配第六代 NVLink 和 CX9 网卡。预计 FP4 精度的推理算力将达到 3.6ExaFLOPS，FP8 精度的训练算力将达到 1.2ExaFLOPS，性能是 Blackwell Ultra NVL72 的 3.3 倍。同时还配备了 HBM4，带宽为 13TB/s；拥有 75 TB 的快速内存，容量是前代的 1.6 倍；支持 NVLink 6，带宽为 260 TB/s，是前代的 2 倍。支持 CX9 网卡，总带宽为 28.8 TB/s，是前代的 2 倍；

2027 将上市 Rubin Ultra NVL576 机柜：FP4 精度的推理和 FP8 精度的训练算力分别是 15ExaFLOPS 和 5ExaFLOPS，14 倍于 Blackwell Ultra NVL72。

图4：英伟达 Rubin 系统-预计将于 2026 年亮相



资料来源：Nvidia，中国银河证券研究院

图5：Vera Rubin Ultra NVL576-预计将于 2027 年亮相



资料来源：Nvidia，中国银河证券研究院

（二）1.6T 硅光引擎 CPO 交换机推出，光电器件有望 2H25 量产

NVIDIA 于 3 月 19 日正式首次亮相了基于 1.6T 硅光引擎的 CPO(共封装光学)交换机系列，包括 Spectrum-x 和 Quantum-x CPO 交换机。黄仁勋表示，NVIDIA 光子学团队携手产业伙伴推动 CPO 联合创新，包括首款采用微环调制器 (MRM) 1.6T 硅光 CPO 芯片，首个采用 TSMC 制程 3D 堆叠硅光子引擎，搭载了高输出功率激光器和直连光纤连接器。Spectrum-x 硅光 CPO 将于 2025 年下半年商用，Quantum-x 硅光 CPO 预计 2026 年下半年商用。值得一提，微环调制器(MRM)和马赫曾德尔调制器 (MZ) 是 CPO 两大调制器技术，除了英伟达采用 MRM，CPO 技术引领者博通也在 TSMC 合作，在 3nm 工艺调试成功 CPO MRM，预计 2025 年初可以交付样品，有望在 2025 年下半年量产 1.6Tbps 光电器件。

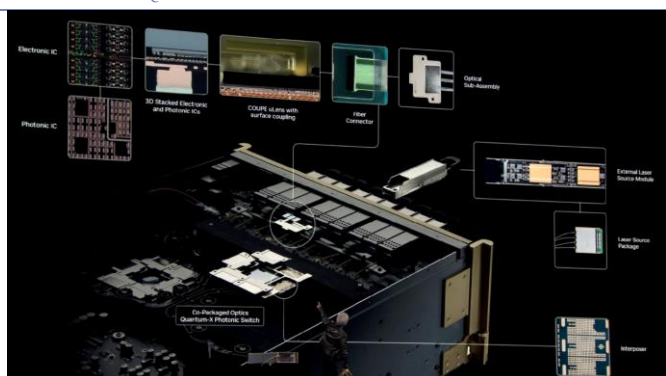
Quantum-x 硅光 CPO：在 115.2Tb/s Quantum-x 光交换机共有 2 个 CPO 模组，一个封装模组有 Quantum-X800 ASIC、6 个光学组件合计 18 个硅光引擎构成，Quantum-X800 ASIC 具备 28.8Tb/s 吞吐量，基于 TSMC 4N 工艺达到 1070 亿个晶体管。CPO 模组单个直连光学组件含有 3 个基于 Interposer 中阶层的硅光引擎(合计 18 个光引擎)和 3 个小型插拔式连接器，可完成 4.8Tb/s 吞吐量。每个硅光引擎皆采用 200Gb/s 微环调制器，可节省 3.5 倍功耗。硅光引擎中的光电芯片以 3D 堆叠集成在衬底上，通过光纤阵列将光信号向外部输出。此外，硅光引擎采用了 TSMC N6 工艺 2.2 亿个晶体管，单片集成 1000 个光子器件。

外部连接方面：Quantum-x 光交换机端口采用 1152 单模光纤 MPO 连接器。外置光源 (ELS) 反面，单个外置光源搭载 8 个激光器 (即 200 毫瓦 CW-DFB)，具有自动温度追踪，波长和功率稳

定性。整体上，一台 Quantum-x 光交换机搭载 2 个 CPO 模组、18 个外置光源、144 根 MPO 连接器，合计 4460 亿个晶体管和 115Tb/s 吞吐量。

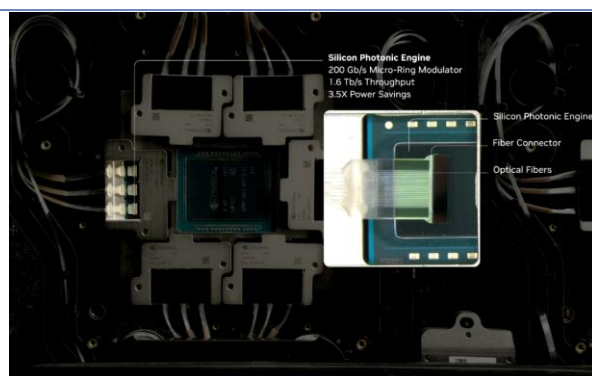
根据讯石光通信援引英伟达口径，基于 NVIDIA CPO 共封装光学平台，NVIDIA 打造世界最先进的 CPO 光交换机系统，包括 128 端口 800G 和 512 端口 800G 的 Spectrum-x 系列 409.6Tb/s 与 102.4Tb/s 光交换机，以及 144 端口 800G Quantum-x 115.2Tb/s 光交换机，预计 2025 年下半年上市。

图6：英伟达 Quantum-X CPO 交换机系统-CPO 部分



资料来源：Nvidia，中国银河证券研究院

图7：硅光引擎组件



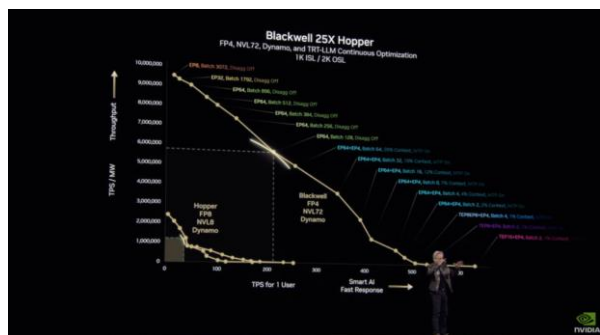
资料来源：Nvidia，讯石光通信，中国银河证券研究院

（三）Nvidia Dynamo 及 Llama Nemotron 发布，软件层面持续发力

在软件领域，Nvidia 宣布了 Nvidia Dynamo——一个专注于简化推理部署和扩展的开放 AI 引擎栈。Dynamo 是一个专为推理、训练和跨整个数据中心加速而构建的开源软件，在现有 Hopper 架构上，Dynamo 可让标准 Llama 模型性能翻倍。而对于 DeepSeek 等专门的推理模型，其智能推理优化还能将每个 GPU 生成的 token 数量提升 30 倍以上。Semianalysis 认为其有可能颠覆 VLLM 和 SGLang，提供 VLLM 所不具备的多种功能，并以更高的性能运行。结合硬件层面的创新，Dynamo 将再次推动吞吐量与交互性曲线的右移——特别是在高交互性用例中提升吞吐量。芯智讯认为，Dynamo 与 CUDA 作为 GPU 编程的底层基础不同，Dynamo 是一个更高层次的系统，专注于大规模推理负载的智能分配和管理。它负责推理优化的分布式调度层，位于应用程序和底层计算基础设施之间。但就像 CUDA 十多年前彻底改变了 GPU 计算格局，Dynamo 也可能成功开创推理软硬件效率的新范式。

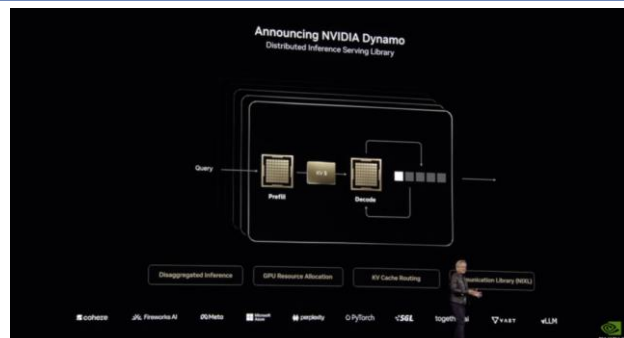
根据 SemiAnalysis，Dynamo 带来了当前推理栈中不具备的几项关键功能：智能路由器（Smart Router）、GPU 规划器（GPU Planner）、改进的 NCCL 集体通信（Improved NCCL Collective for Inference）、NIXL——NVIDIA 推理传输引擎（NIXL - NVIDIA Inference Transfer Engine）、NVMe KV 缓存卸载管理器（NVMe KV-Cache Offload Manager）。

图8：随着 Dynamo 的加持，Blackwell 对 Hooper 有巨大优势



资料来源：Nvidia，中国银河证券研究院

图9：Dynamo 将 LLM 不同阶段分配至不同 GPU 中



资料来源：Nvidia，中国银河证券研究院

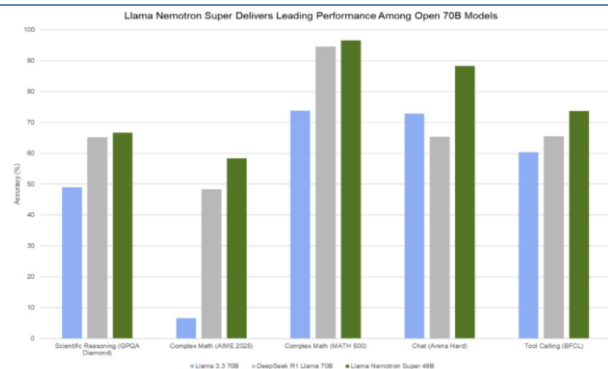
模型方面，英伟达发布 NVIDIA Llama Nemotron。本次 GTC 2025，英伟达发布由 Llama 模型衍生出的训练模型 Nvidia Llama Nemotron，主打高效、准确。相较于 Llama，该款模型经过算法修剪优化，不仅更为轻量级（48B），还拥有类似 o1 的推理能力。与 Claude 3.7 和 Grok 3 一样，Llama Nemotron 模型内置了推理能力开关，用户可选择是否开启。这个系列分为三档：入门级的 Nano、中端的 Super 和旗舰 Ultra，每一款都针对不同规模的企业需求。

图10: NVIDIA Llama Nemotron 模型



资料来源: Nvidia, 中国银河证券研究院

图11: 推理效率对比-Llama Nemotron 及同类竞品



资料来源: 芯智讯, 中国银河证券研究院

随着 NVIDIA Llama Nemotron 的推出，英伟达 AI Agent 时代有望到来。OpenAI、Claude 等大厂逐步通过 DeepResearch、MCP 建立起了 Agent 的基础，英伟达有望以 Llama Nemotron 推理模型作为基础，通过 NVIDIA AIQ 架构作为 Agent 模板进行相关方向探索。

图12: 英伟达 Servicenow 系统



资料来源: Nvidia, 中国银河证券研究院

二、投资建议

通过此次英伟达 GTC 2025 进行分析，我们可以发现当前算力板块在推理应用的不断发展中并非呈现出算力需求下降的态势，而是算力需求得到进一步的刺激，当前市场普遍认为推理应用的不断增长将挤压 2026 年甚至更长时间段内算力板块的成长空间，我们认为该种判断可能同算力行业的实际情况有所偏离，在算力需求总体增长的情况下，我们认为当前算力相关板块仍具备较大投资价值，建议优选空间较大且动能充足的细分子板块：

运营商：低估值高成长，国家政策支持下算力新业务发展有望超预期。

运营商盈利能力、现金流资产不断改善、资产价值优势凸显，持续增加分红回馈股东，相对历史估值和国外水平，通信运营商均处于估值低位。总体来说，运营商业绩持续增长或超预期，5G“收获期”大有可为。当前运营商云业务发展如火如荼，DeepSeek 对于成本端的降低有望协同运营商云业务部署以及运营商的海量数据资产，推动运营商第二曲线的快速增长。

建议关注：中国移动、中国联通、中国电信等。

光通信：技术迭代产品量价齐升，技术壁垒增强竞争格局有望边际改善。

AIGC 引领新一轮科技革命，以 DeepSeek 为首的推理模型对于成本端的降低或将推动应用端的繁荣，继而反哺推理侧模型的快速迭代，在本次 GTC 大会中也有所体现，有望推动应用端的进一步发展，带动光器件需求量持续高增。光模块 100G/200G→400G→800G→1/6T 迭代速率持续提升，带来产品量价齐升有望延续，带来业绩高增持续可期。同时国内的算力部署有望推动国内光模块产业链景气度提升，带来较强的规模效应。

建议关注：中际旭创、新易盛、天孚通信、光迅科技、华工科技等。

光芯片产业链国产化进程加速，复杂国际经济形势下渗透率有望进一步提升。

以 DeepSeek 为首的推理算力对于成本端及训练精度的降低或将使得推理侧对光芯片的技术需求边际改变，国产光芯片在推理侧算力部署中具备成本优势且技术可靠性较强，产业链渗透率有望跟随推理侧算力部署的增加而有所上升，在复杂国际形势下，国产光芯片在推理侧部署的可靠性及成本优势有望将进一步提升自身在采购链条中的话语权和份额。

建议关注：源杰科技、仕佳光子等。

以 DeepSeek 为首的推理算力增长推动 AIDC 行业加速发展，供需改善。

DeepSeek 的推出，对于算力的需求并非缩减而是持续增加，当前算力规模并不足以支撑 DeepSeek 全时间段大规模运行，算力的强需求将带动现有 AIDC 的上架率提升，以及新机房建设需求也成为应用端发展的确定性方向。

建议关注：AIDC 相关光环新网、数据港等。

三、风险提示

AI 发展不及预期的风险：若 AI 发展不及预期，则算力需求结构可能有所变化，若 AI 发展不及预期，则算力发展节奏可能有所减缓，带动相关板块增速下降；

全球地缘政治不确定性的风险：当前全球地缘政治具备一定的不确定性，若部分地区贸易保护力度增加或关税增长，则算力相关板块即便在需求无虞的情况下，对利润的弹性也将边际减弱；

算力发展不及预期的风险：算力板块受硬件发展速度以及建设意愿等因素影响较大，若算力发展不及预期，则软件发展也将有所减缓。

图表目录

图 1: Blackwell Ultra NVL72-AI 推理为主4

图 2: Blackwell Ultra NVL72 同上代产品对比4

图 3: 英伟达 RTX 与 DGX 产品矩阵融合, 覆盖场景进一步加强4

图 4: 英伟达 Rubin 系统-预计将于 2026 年亮相5

图 5: Vera Rubvin Ultra NVL576-预计将于 2027 年亮相5

图 6: 英伟达 Quantum-X CPO 交换机系统-CPO 部分6

图 7: 硅光引擎组件.....6

图 8: 随着 Dynamo 的加持, Blackwell 对 Hooper 有巨大优势6

图 9: Dynamo 将 LLM 不同阶段分配至不同 GPU 中6

图 10: NVIDIA Llama Nemotron 模型.....7

图 11: 推理效率对比-Llama Nemotron 及同类竞品7

图 12: 英伟达 Servicenow 系统7

表 1: NVLink 历代性能参数对比及进化3

分析师承诺及简介

本人承诺以勤勉的执业态度，独立、客观地出具本报告，本报告清晰准确地反映本人的研究观点。本人薪酬的任何部分过去不曾与、现在不与、未来也将不会与本报告的具体推荐或观点直接或间接相关。

赵良毕，通信&中小盘首席分析师，科技组组长。北京邮电大学通信硕士，复合学科背景，2022 年加入中国银河证券。8 年中国移动通信产业研究经验，6 年证券从业经验。曾获得 2018/2019 年（机构投资者 II-财新）通信行业最佳分析师前三名，2020 年获得 Wind（万得）金牌通信分析师前五名，获得 2022 年 Choice（东方财富网）通信行业最佳分析师前三名。

免责声明

本报告由中国银河证券股份有限公司（以下简称银河证券）向其客户提供。银河证券无需因接收人收到本报告而视其为客户。若您并非银河证券客户中的专业投资者，为保证服务质量、控制投资风险、应首先联系银河证券机构销售部门或客户经理，完成投资者适当性匹配，并充分了解该项服务的性质、特点、使用的注意事项以及若不当使用可能带来的风险或损失。

本报告所载的全部内容只提供给客户做参考之用，并不构成对客户投资咨询建议，并非作为买卖、认购证券或其它金融工具的邀请或保证。客户不应单纯依靠本报告而取代自我独立判断。银河证券认为本报告资料来源是可靠的，所载内容及观点客观公正，但不担保其准确性或完整性。本报告所载内容反映的是银河证券在最初发表本报告日期当日的判断，银河证券可发出其它与本报告所载内容不一致或有不同结论的报告，但银河证券没有义务和责任去及时更新本报告涉及的内容并通知客户。银河证券不对因客户使用本报告而导致的损失负任何责任。

本报告可能附带其它网站的地址或超级链接，对于可能涉及的银河证券网站以外的地址或超级链接，银河证券不对其内容负责。链接网站的内容不构成本报告的任何部分，客户需自行承担浏览这些网站的费用或风险。

银河证券在法律允许的情况下可参与、投资或持有本报告涉及的证券或进行证券交易，或向本报告涉及的公司提供或争取提供包括投资银行业务在内的服务或业务支持。银河证券可能与本报告涉及的公司之间存在业务关系，并无需事先或在获得业务关系后通知客户。

银河证券已具备中国证监会批复的证券投资咨询业务资格。除非另有说明，所有本报告的版权属于银河证券。未经银河证券书面授权许可，任何机构或个人不得以任何形式转发、转载、翻版或传播本报告。特提醒公众投资者慎重使用未经授权刊载或者转发的本公司证券研究报告。

本报告版权归银河证券所有并保留最终解释权。

评级标准

评级标准	评级	说明
评级标准为报告发布日后的 6 到 12 个月行业指数（或公司股价）相对市场表现，其中：A 股市场以沪深 300 指数为基准，新三板市场以三板成指（针对协议转让标的）或三板做市指数（针对做市转让标的）为基准，北交所市场以北证 50 指数为基准，香港市场以恒生指数为基准。	行业评级	推荐：相对基准指数涨幅 10%以上
		中性：相对基准指数涨幅在-5%~10%之间
		回避：相对基准指数跌幅 5%以上
	公司评级	推荐：相对基准指数涨幅 20%以上
		谨慎推荐：相对基准指数涨幅在 5%~20%之间
		中性：相对基准指数涨幅在-5%~5%之间
		回避：相对基准指数跌幅 5%以上

联系

中国银河证券股份有限公司 研究院

深圳市福田区金田路 3088 号中洲大厦 20 层

上海浦东新区富城路 99 号震旦大厦 31 层

北京市丰台区西营街 8 号院 1 号楼青海金融大厦

公司网址：www.chinastock.com.cn

机构请致电：

深广地区：程 曦 0755-83471683 chengxi_yj@chinastock.com.cn
 苏一耘 0755-83479312 suyiyun_yj@chinastock.com.cn
 上海地区：陆韵如 021-60387901 luyunru_yj@chinastock.com.cn
 李洋洋 021-20252671 liyangyang_yj@chinastock.com.cn
 北京地区：田 薇 010-80927721 tianwei@chinastock.com.cn
 褚 颖 010-80927755 chuying_yj@chinastock.com.cn