

超节点设计实践

王小泉

远图产品总工



目录

CONTENTS

1 Scale up带宽

2 HBD Size和可扩展性

3 互联拓扑和协议

4 GPU功耗与散热

5 系统架构设计

6 信号完整性评估

7 供电方案评估

8 系统管理方案

9 超节点发展趋势

01 单GPU支持的最大Scale up带宽



超高带宽

Scale up

7.2Tbps

VS

Scale Out

800Gbps

Blackwell单GPU总带宽

核心实践：定义互联的物理上限

明确单GPU芯片间点对点互联的最大带宽能力，通常以“Serdes速率Gbps × Lane数”表示。如112G × 64 Lane，即单通道速率112Gbps，共64个通信通道，以此确定对外的连接数量和连接器选型

未来趋势：速率、通道的演进

GPU高速Serdes速率与通道规模持续扩容：

当前主流为112G Serdes，下一阶段将全面转向224G Serdes。

通道数量从32Lane向96Lane甚至128Lane扩容。单通道速率与通道数量的双重提升，将成倍拓展GPU间点对点互联总带宽

448G Serdes受半导体工艺、信号完整性、散热功耗、PCB板材多重技术约束，产业链成熟度不足，短期无法实现大规模商用落地，仍需长期技术攻关

02 单层HBD Size和可扩展性

核心实践：定义与关键参数

定义：HBD(High Bandwidth Domain)

一组通过超高速 Scale-Up 专用互联链路紧密耦合的 XPU集合，域内算力芯片之间的点对点通信带宽、时延显著优于跨域（HBD 之间）的通用以太网 Scale-Out 网络。

主流规格

行业常见的配置为32、64、128卡等，需根据实际算力需求与组网成本进行选型匹配

可扩展性

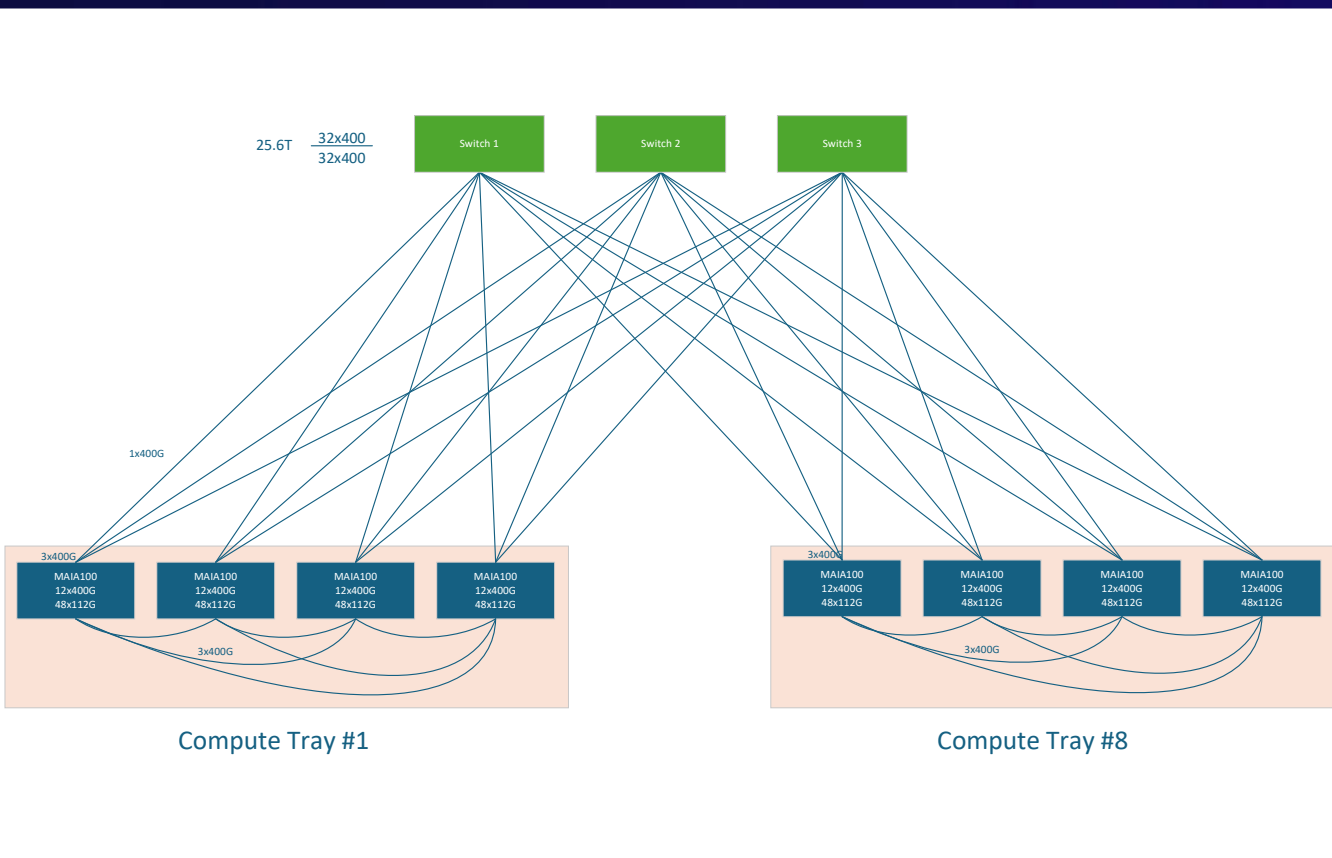
直接决定超节点内部Switch单元是否需要对外Scaleup扩展接口
NVL36: 可扩展
NVL72: 不可扩展



未来趋势：机柜扩容与算力密度提升

未来超节点机柜将突破600mm标准宽度，采用900mm/1200mm双宽规格和定制化宽度机柜。
单机柜GPU集成密度将从64/72卡升级至128/144卡，并向256卡甚至512卡迈进。

03 互联拓扑和互联协议



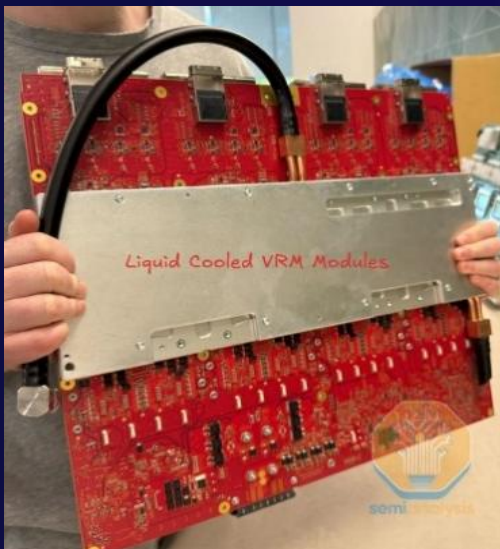
核心实践：拓扑与协议的选型平衡

优先选择CLOS、Mesh、或者混合拓扑保障高带宽低时延；
协议兼顾以太网通用性与厂商专用链路性能，平衡成本、扩展性与运维复杂度。

未来趋势：高速互联协议由封闭走向开放

打破NVLink等封闭协议壁垒，华为UB总线、海光HSL协议等国产自主协议与ESUN/SUE、UALink等通用开放协议并行发展，推动算力基础设施自主可控、多元共存。

04 单GPU最大功耗及散热需求



未来趋势：散热与功耗双重升级

- 单芯片功耗全面升级，垂直供电成为主流：**单芯片功耗突破传统数百瓦区间，稳步向3KW至5KW超高功耗区间演进。64卡1KW整机柜功耗约100KW；128卡3.5KW整机柜功耗突破600KW。
- 整机柜功耗量级跨越，迈入兆瓦级区间：**未来超节点整机柜峰值功耗将突破1MW，倒逼机房配电、冷却、消防体系同步重构。
- 散热体系全面升级，全液冷成为主流：**风冷和风液混合将逐步淘汰，下一代超高密超节点将全面采用全液冷架构（冷板式/浸没式），将数据中心PUE控制在1.1以内。

• 单GPU满负载功耗 (TDP)

明确单颗GPU在算力满负载运行下的最大热设计功耗，这是规划供电容量与散热方案的核心基础数据

• 分级决策：风冷向液冷技术迭代

根据功耗等级阶梯式选择散热方案：10-40KW中低功耗-高效风冷，单柜>50KW高功耗场景-风液混合散热，高功率芯片-冷板散热，确保热量被快速带走，维持芯片运行在最佳温区

• 全机柜级全液冷必要性研判

针对超高密度部署场景，需评估对整个机柜或节点实施全液冷散热，从根本上解决传统风冷在散热效率、噪音和能耗上的瓶颈问题

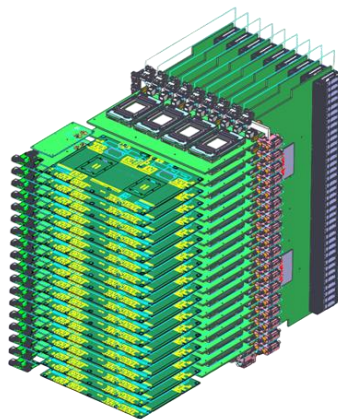
05 系统架构评估



Cable Tray (线缆互联)

优点：结构简单、维护方便、成本低

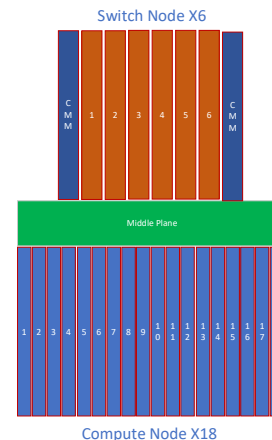
缺点：机柜布线复杂、占空间、散热受限，不利于高密度算力集群部署扩展



直接正交 (OD) 架构

优点：取消中间线缆，信号路径短、时延低、传输效率高，维护便捷

缺点：设计精度要求高，装配工艺复杂，硬件与制造成本显著增加



中背板 (MiddlePlane) 方案

优点：通过背板集中走线，布线极简整洁，系统可靠性与稳定性极强

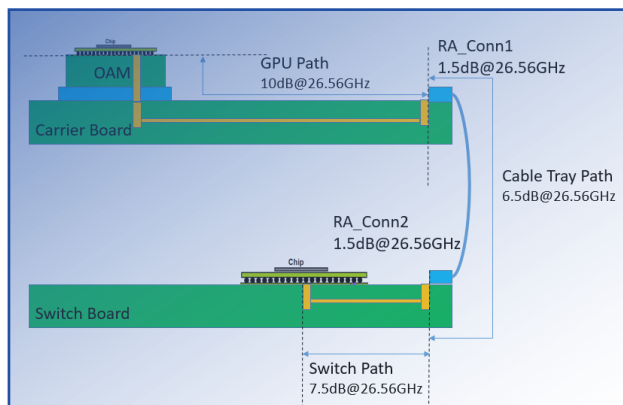
缺点：硬件成本高昂，背板接口固定，扩容与配置调整灵活性弱，难以适配快速变化的算力需求。

未来趋势：三种架构长期共存

无统一“最优架构”，厂商将依据成本、带宽、可靠性需求差异化取舍，三种方案长期并行，适配不同算力场景

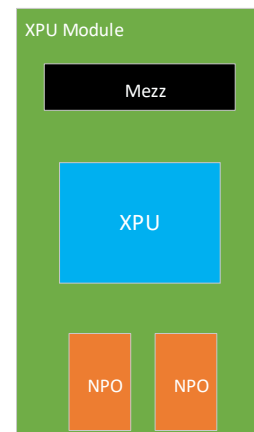
06 信号完整性评估

核心实践：SI评估与Retimer决策



随着高速信号传输速率提升，信号衰减、串扰与抖动问题凸显。Retimer芯片可重新定时、放大信号，是补偿损耗、延长传输距离的关键组件。

未来趋势：光电分层互联成为标准



机柜内部短距离互联保留铜缆，跨机柜、跨机房远距离大规模集群组网统一切换光互联。

GPU侧搭载NPO（近封装光学）模块成为CPO/OIO落地前的中长期过渡方案。

SI 信号评估

全面分析衰减、串扰与抖动，量化信号完整性指标

Retimer 补偿

重定时放大信号，有效抵消链路损耗，提升传输稳定性。功耗/成本考量

场景化决策

依据仿真结果，在长距或高密度链路中按需部署组件

分层部署架构

机柜内短距用铜缆，跨机柜、跨机房长距统一采用光互联

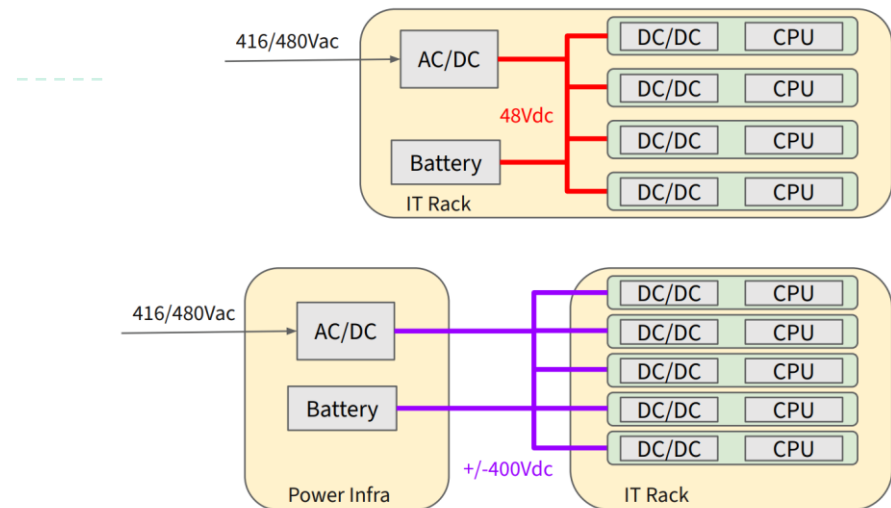
NPO 过渡方案

GPU侧搭载近封装光学模块，作为CPO落地前的中长期选择

向 CPO 演进

待技术成熟后，逐步平滑过渡至性能极致的共封装光学架构

07 系统供电评估



多元供电架构

灵活选择机柜级供电方案，适配传统48V分布式供电或新兴高压直流母线架构，夯实系统底座

技术选型与考量

评估54V/800V Busbar及SideCar等技术，确保方案满足高功率、高效率、高可靠性与空间优化需求

未来趋势

• 机柜供电架构高压化

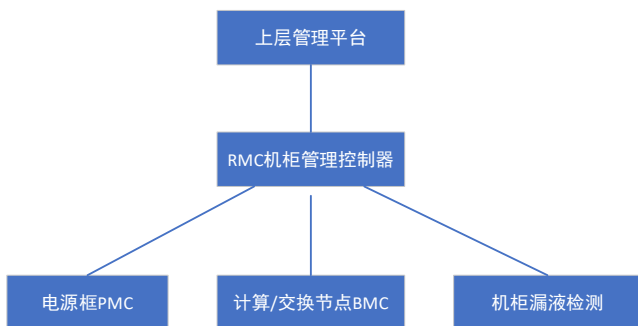
传统48V低压母线将被±400V/800V高压直流HVDC全面替换，以降低电流、减少损耗、节省空间，匹配兆瓦级机柜需求

• 整机柜功耗量级跨越

超节点整机柜峰值功耗演进：50KW-100KW-200KW-500KW-1MW, 正式进入兆瓦级时代。

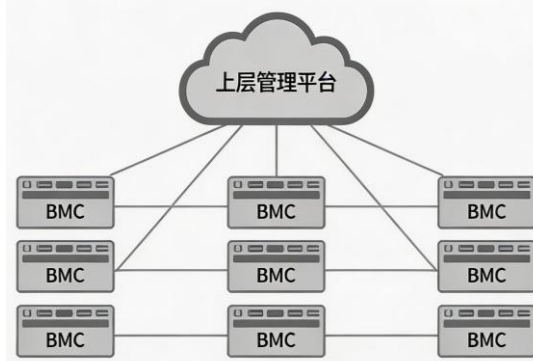
08 超节点系统管理方案

集中式管理 (RMC)



以机柜管理控制器(RMC)为核心中枢, 将整机柜视为单一逻辑资源池。通过硬件解耦技术, 实现计算、存储、网络资源的灵活组合与统一调度, 构建一体化的基础设施管控体系。

分布式管理 (BMC)



基于节点自带的BMC实现自治管理。各节点独立运行, 通过标准网络协议接入上层管理平台, 支持多品牌、多型号硬件的异构混合部署。

核心特点

实现机柜级统一监控与控制, 资源全局池化, 通过单点入口大幅简化大规模集群的运维复杂度。

核心优势

支持功率封顶、漏液检测等精细化管控; 支持批量固件升级与统一审计, 显著提升资源利用率

潜在局限

RMC故障可能导致整机柜管理瘫痪; 软硬件深度绑定, 存在较强的厂商锁定效应, 灵活性不足

核心特点

节点自治运行, 无全局单点故障风险; 兼容主流硬件生态, 具备极强的异构兼容性与部署灵活性。

核心优势

去中心化高可用, 支持近乎无限的水平弹性扩展

潜在局限

缺乏机柜级全局视角, 批量操作效率较低;

09 超节点发展趋势

1. 更高算力密度

通过架构革新与工艺升级，实现算力单元的极致压缩，最大化单位空间内的计算性能，夯实超节点硬件底座。

2. 更高互联带宽

构建高速互联网络，实现节点间低时延、高吞吐的数据传输，消除数据流动瓶颈，保障大规模集群协同效率

3. 更高能效供电散热

突破传统供电与散热技术瓶颈，以更优能耗比支持高密度算力运行，确保系统长期稳定、绿色可持续

4. 开放自主互联生态

打造开放兼容、自主可控的技术体系，推动软硬件标准化与模块化，汇聚产业合力共建繁荣的超节点创新生态

5. 异构融合算力

融合通用CPU+GPU+NPU+LPU 等多元算力架构，实现优势互补，为复杂AI任务提供最优的算力支撑方案

6. 全域智能软件调度

依托AI算法实现算力资源的全局感知与动态调度，优化任务分配与负载均衡，让每一份算力都发挥最大价值

Thank You!