

超节点发展报告



联合编写单位



协作单位

中国电子技术标准化研究院、GCC 全球计算联盟、国家信息中心

CONTENTS

目录

01	●	前言	06
02	●	大模型对基础设施的挑战	07
		2.1 通往通用人工智能之路：最新大模型发展动态	07
		2.2 AI 技术从单点能力突破迈向系统能力创新	07
		2.3 大模型计算基础设施的挑战	09
		2.4 小结	11
03	●	超节点的出现与演进	12
		3.1 全球产业的演进路线：从硬件聚合到系统构建	12
		3.2 超节点技术产业生态发展格局	13
		3.3 小结	13
04	●	超节点基础定义与特征	14
		4.1 技术特征	14
		4.1.1 基础特征：大带宽、低时延、内存统一编址	14
		4.1.2 扩展特征：多级缓存池化、资源灵活配比	15
		4.2 系统特征	16
		4.2.1 超大规模	16
		4.2.2 超高可靠	18
		4.2.3 灵活切分	20
05	●	超节点应用案例	21
		5.1 支撑大模型创新及云服务场景	21
		5.2 加速人工智能科学计算，服务算法创新	22
		5.3 助力行业企业智能化升级	22
06	●	总结和展望：迈向未来计算的下一个十年	24
07	●	参考文献	26

序言 1

当我们站在人工智能大模型技术飞速发展的十字路口，一个清晰的趋势已然浮现：大模型正沿着“规模定律”不断演进，从预训练扩展到覆盖预训练、后训练、逻辑推理的全流程，其参数与集群规模实现“双万”跨越，行业模型落地需求专业化。

传统的服务器集群架构在这场变革中瓶颈愈发明显。千亿级模型一次梯度同步产生的 TB 级数据让传统以太网带宽难以承受；同时，伴随算力规模扩大，万级处理器带来的故障常态化，对自动化运维与 RAS 能力提出了更高要求。在这样的背景下，超节点的出现成为了面向大模型未来发展的必然趋势。

超节点并非简单的硬件堆砌，它的实现离不开基础技术、系统能力与可落地性的三方协同。基础技术是超节点的根基，其具备超高带宽互联、内存统一编址等技术特征，通过近乎无阻塞的高带宽互联，将数百上千个 AI 处理器编织为一个逻辑统一的高密度计算体，为高效计算提供了底层支撑。系统能力则是超节点高效运转的保障，它需要具备大规模、高可靠、多场景等系统特征。大规模的组网能力突破了单机扩展的硬件限制，为大规模算力聚合提供架构支撑；高可靠的运行特性化解了网络、计算、存储等子系统的故障风险，保障集群作业的连续性；多场景的适配能力则能通过精细化资源调度等机制，满足不同业务需求，最大化释放算力价值。

本文系统性地提出并论证了“超节点将成为 AI 时代的核心计算单元”这一重要观点，清晰地呈现了超节点的基础定义与特征，包括技术层面的基础特征和扩展特征，以及系统层面的大规模、高可靠、多场景特征。同时，通过分析全球产业的演进路线、超节点稳定性的核心挑战以及技术产业生态发展格局，为产业界指明了超节点的发展方向。

在未来计算的下一个十年，超节点无疑将成为推动 AI 技术发展的关键力量。这份发展报告为我们提供了宝贵的理论指导和实践参考，相信在产业界的共同努力下，超节点将不断成熟完善，为人工智能的持续突破和广泛应用奠定坚实的算力基础。

中国工程院院士、清华大学计算机系教授

郑纬民

序言 2

大模型正以不可逆转之势为全球计算领域带来跨越式变革。从生成式 AI 到 Agentic AI 再到 Physical AI，大模型持续提升解决复杂问题的能力，并向物理世界延伸。大模型技术及能力演进，驱动 AI 系统负载变化，需要一套系统架构满足未来发展需求，超节点成为 AI 基础建设的共识。

超节点架构引领技术革新，重构计算能力边界。超节点架构依托高速互联技术，将大带宽的互联范围，从单台服务器扩展到整机柜以及跨机柜的大规模集群，超节点域内可达百 GB/s 级通信带宽、纳秒级时延、TB 级超大内存，实现集群能力跃迁。相较“服务器集群”，超节点代表的是弹性、池化、开放的系统能力：既能以极致吞吐支撑万亿参数训练，也能以低时延满足企业级大规模推理的刚性需求。

昇腾 AI 坚持架构创新，开源开放，共建产业生态。昇腾 AI 经过 6 年快速发展，已成长为中国 AI 算力第二平面的坚实基础，并通过软硬件开源开放，建立生态兼容、共建共享的昇腾 AI 生态。在基础硬件层面，昇腾持续引领技术架构，打造领先产品，实现业界最大规模的 384 超节点产品，并在下一代将扩展至 8192，持续领先；在基础软件层面，通过一套架构满足不同代际产品的持续演进，同时秉承开源开放的策略，将核心计算架构 CANN、Mind 系列应用使能软件全面开源开放，同时结合对 PyTorch 等主流框架的全面兼容和体系化工具链，旨在最大限度地降低开发门槛，加速开发者和社区的融入。如今，昇腾 AI 的算力底座已支撑起互联网、金融、政务、制造等数十个行业的智能化转型，累计服务超过 10 万家企业客户。

携手生态伙伴，共筑产业 AI 生产力。面向 AI 产业的广阔前景，我们将以技术创新为本，构建持续领先的产品及解决方案，并将能力开放出来，支持伙伴打造多样化产品，并为企业提供有竞争力的解决方案，覆盖越来越多的行业场景。我们将与客户、伙伴形成紧密协同的价值共同体，加速产业界的智能化升级与创新，将人工智能带入丰富的行业场景，让智能无所不及。

华为公司董事、ICT BG CEO

杨超斌

序言 3

我们正站在一个智能变革涌动的时代潮头。以大模型为代表的人工智能技术，成为驱动千行百业颠覆性变革的核心力量。大模型所展现出的涌现能力与通用潜能，正在重构人类对创新的想象边界，但同时也对底层智算基础设施提出更高要求和挑战：模型参数规模从千亿迈向万亿乃至更高，训练数据量呈指数级增长，传统松散耦合的集群架构已难以满足高效的计算需求，智算基础设施正开始新一轮的技术革新。

在此背景下，超节点应运而生。它并非偶然的技术产品迭代，而是智算需求与系统创新深度共振的结果，具有划时代的重要意义。超节点超越简单的硬件集成，代表一种全新的构建哲学：以系统化、一体化的设计思维，将计算、存储、网络与运维管理深度融合，锻造出高性能、高效率、高可靠的单一逻辑实体。它标志着一个全新时代的开启——智算基础设施正从松散组合的算力堆叠阶段，迈入软硬协同、全局优化的超节点阶段，旨在有效破解超大规模 AI 训练与推理中所面临的扩展性瓶颈、效率损耗与能耗墙难题，为 AI 的持续创新提供坚实、高效、绿色的算力基座。

为系统分析超节点技术的发展逻辑、技术创新、产业价值以及未来趋势，我院与华为及相关单位共同开展研究，编制《超节点发展报告》。报告以“需求—技术—应用—展望”为主线，从大模型对智算基础设施的机遇与挑战入手，深入剖析超节点技术的发展动因，对超节点技术的发展历程及技术特征进行梳理，为各行业的应用落地提供参考。

我们坚信，超节点是未来构建高效可靠算力优势的关键抓手。超节点的成熟与普及，对于提升整体智算能力、促进 AI 赋能千行百业具有至关重要的意义。我们期待本报告能凝聚产业共识，推动超节点技术从“技术探索”走向“落地应用”，加速我国智算基础设施发展，为全球智算产业创新贡献中国智慧。

中国信息通信研究院副院长

魏亮

推荐语

人工智能高速演进背景下，算力需求呈指数级增长，大模型竞争已进入“参数规模摸高”与“训练效率提升”并行的新阶段。Scaling Law（规模定律）将以多元形态长期生效，持续推动人工智能技术突破能力边界，而超大规模 Transformer、MoE（混合专家模型）、稀疏注意力模型等，已成为可扩展模型的核心架构方向。在复杂的混合并行策略下，随着并行规模持续扩大，系统节点间通信带宽与可用显存容量成为制约大模型可扩展性的瓶颈，亟需计算架构创新以满足未来更大规模模型训练的需求。超节点架构突破传统互联瓶颈与共享协议限制，不断突破系统性能上限，成为多样化算力集群技术未来演进的必然趋势。本发展报告系统梳理了超节点技术架构的国内外演进路径与生态发展格局，清晰界定了超节点需具备的技术特征与系统属性，为产业界提供了具有前瞻性的洞见和系统标准参考，助力行业在算力发展中找准技术方向，推动算力从规模堆砌走向效率跃迁。

国家超级计算广州中心主任 卢宇彤

当前，千亿乃至万亿参数的大模型与 MoE 等先进架构的兴起，对计算基础设施提出了前所未有的苛刻要求。传统的硬件堆叠模式已难以满足其对于算力规模、通信效率及系统稳定性的需求。《超节点发展报告》深刻阐述了，必须从单纯的硬件聚合，迈向以“系统工程”思想为核心的创新构建。超节点通过超高带宽互联、内存统一编址等关键技术，实现了计算、存储、网络资源的深度融合与高效协同，其大规模灵活组网与高可靠运行的系统优势，是构建稳定、高效、易用的新一代算力系统的必然路径。超节点是支撑未来复杂 AI 计算任务的关键基石，本发展报告对其技术内涵与应用价值的系统梳理，对产业生态发展具有重要的指导意义。

中国电子技术标准化研究院 副院长 范科峰

在大模型飞速发展与应用需求爆发的时代，AI 基础设施面临诸多挑战，传统的计算架构已难以满足高效、大规模的训练和推理需求。《超节点发展报告》深入剖析了超节点如何凭借创新架构，构建高效协同机制，让算力、算法、数据得以深度融合，打破单点能力的局限，实现“系统能力”的创新，提升大模型训练的效率，显著降低推理时延。《超节点发展报告》为业界如何把握这一前沿趋势、共同推进全球 AI 的技术进步提供了重要参考。

GCC 全球计算联盟理事长 华中科技大学教授 金海

世界正进入一个对 AI 算力需求空前增长的时代，大模型训练成本的指数级增长，正迫使行业将重心从单纯的算力堆砌转向底层计算架构的根本性创新。每一次计算范式的更替，都会重塑产业版图。今天，生成式人工智能正把基础设施从“堆叠服务器的集群”，推向“像一台巨型计算机运作的集成单元”。这不是简单的规模扩张，而是一场关于带宽、能效与系统韧性的全面再造。预计到 2030 年，相关基础设施投资将接近 7 万亿美元（Noffsinger et al., 2025）^[1]。计算范式正从通用数据中心转向专为 AI 设计的“超节点”，这正在彻底改变数字基础设施的经济模型与设计理念：资本开支、能耗曲线、网络与内存比重、运维与可靠性能力，都会被重新定价与重构。

本报告提出并论证：“超节点”将成为 AI 时代的核心计算单元。它通过近乎无阻塞的高带宽互联，把数十到数百乃至数千个 AI 处理器（本文中提及的 AI 处理器泛指用于人工智能计算的加速器，如神经网络处理器（NPU）和图形处理器（GPU）等）编织为一个逻辑统一的高密度计算体；通过软硬件协同和智能编排，把训练与推理的双重诉求在同一平台上高效切换；通过液冷、供电与可观测性 /RAS 体系，把能效与可用度拉回可持续区间。相较“服务器集群”，超节点代表的是弹性、池化、开放的系统能力：既能以极致吞吐支撑万亿参数训练，也能以低时延满足企业级大规模推理的刚性需求。

我们相信，计算将再次成为增长曲线的起点。当超节点把“算力、带宽、内存、能效与可靠性”整合为一体并可编排时，AI 不只是更强的内容生成器，而是可被工业化复制的认知基础设施。这既是技术路线的抉择，也是产业组织与国家竞争力的选择题。答案取决于我们今天如何设计并投资下一代算力系统。

2.1 通往通用人工智能之路：最新大模型发展动态

AI 大模型正以前所未有的速度发展，行业呈现出模型加速迭代、算力大规模投入和商业化进程加快的特点。全球科技巨头与新兴力量纷纷布局，推动技术边界不断拓展。以 OpenAI 为首的美国公司持续引领潮流，其 GPT-5 模型在多个基准测试中排名第一，采用了能为不同任务匹配最适模型的“Router”架构，并投入数十万 GPU 进行训练与推理。Google 则凭借其强大的生态系统，将 Gemini 2.5 Pro 等自研模型深度整合进搜索、Gmail 等全线产品，并通过包含多项 AI 服务的订阅套餐实现商业价值提升。此外，xAI 的 Grok 4 模型通过投入 20 万 H100 进行后训练，在复杂推理任务上表现卓越，创始人马斯克更强调 AI 与物理世界的交互，计划将其植入特斯拉。

与此同时，中国 AI 力量迅速崛起，推出了一系列性能卓越的大模型。月之暗面 (Moonshot AI) 的 Kimi K2 智能助手在推理、编程和工具调用方面进行了重点升级，旨在高效解决用户的复杂问题。深度求索 (DeepSeek) 发布的 DeepSeek V3.1 具备创新的混合推理架构，能同时支持“思考”与“非思考”两种模式，在提升复杂任务处理能力的同时也优化了响应效率。阿里巴巴通义 (Alibaba Qwen) 的 Qwen3 将预训练数据量提升至近 36 万亿 tokens，并引入混合思维模式以提升智能体 (Agent) 能力，其模型家族支持高达 100 万 tokens 的超长上下文处理。智谱 AI (Zhipu AI) 推出了为智能体任务优化的 GLM-4.5 系列，该模型总参数量达 3550 亿，同样采用混合推理模式，并在工具调用成功率上表现优异。

2.2 AI 技术从单点能力突破迈向系统能力创新

总体来看，无论是“思考”模式的引入、Agent 能力的强化，还是开源社区的繁荣，都标志着 AI 技术正从单点能力突破，迈向更通用、更智能的未来。全球 AI 大模型正朝着更大规模、更高效率、更强自主性的方向迈进，这意味着人工智能大模型的发展已进入一个系统性竞争的新阶段。这不仅定义了技术的前沿，也对底层基础设施提出了前所未有的要求。人工智能大模型对计算基础设施的挑战是系统性的，涵盖了算力、通信、功耗和运维等多个维度。

然而，这些看似分散的挑战，其根源几乎都可以追溯到一个核心的驱动理论——“规模定律”。“规模定律”的提出是 AI 发展的里程碑，揭示了模型性能与参数、数据量、计算投入的关系，促使训练从“赌注”走向“可量化投资” (Kaplan et al., 2020)^[2]；钦奇拉法则 (Chinchilla Law)

(Hoffmann et al., 2022) 进一步要求参数与数据按比例协同扩展，把难题从堆算力转向高带宽、低时延、持续数据供给的均衡系统。这一修正的意义是深远的：它将业界的核心难题从单纯“堆砌算力”，转向了如何构建一个能够支撑海量数据持续、高效供给的均衡系统，即对高带宽、低时延的数据传输能力提出了刚性要求。在这一理论指导下，人工智能大模型正沿着“规模定律”的路径，从单一的预训练环节，扩展为覆盖预训练、后训练、逻辑推理的全流程。这一全流程的扩展，不仅全面提升了模型的智力水平，也顺理成章地将基础设施的挑战从过去的“纯算力”问题，升级为“算力 × 数据供给 × 系统编排”的综合性工程难题。

趋势一：基础模型竞赛——参数与集群规模的“双万”跨越

参数从亿级跃迁至万亿级，训练集群从“万卡”走向“十万卡”。从 GPT-1 (1.17 亿) 到 GPT-3 (1750 亿)，再到 GPT-4 (约 1.8 万亿)，2025 年 Llama-4、Kimi K2、xAI Grok4 等模型将万亿级参数与万卡级集群规模确立为新常态（业界预期 GPT-5 或达 10T 与十万卡集群）。头部玩家以 8 - 12 周节奏推新，工程与规模交替领先。截至 2025 年 7 月，中国 433 款大模型完成备案并上线服务^[3]。

趋势二：行业模型落地——后训练与推理需求的专业化

行业客户在基础模型之上经 SFT + RL 形成“行业 R1”，算力从百卡增至千卡。推理端爆发要求低时延 + 高吞吐并重：既要扛训练峰值，又要以高效率、低成本、多租户隔离稳定供给，推动 GPU/NPU 虚拟化与弹性调度。

OpenAI 范例：OpenAI 坚定地遵循“规模定律”，通过投入更多算力和更大参数来保证其基础模型（如 GPT-4/4.5）的持续领先。同时，它基于这些基础模型，通过后训练技术扩展出面向推理的“o 系列”模型，满足更专业的市场需求。

趋势三：大模型训练成本倍数级增长趋势

根据 Cottier, B., et al. (2024)^[4] 的研究分析，大型语言模型（LLM）的训练成本正呈现出惊人的倍数级增长趋势。前沿模型训练成本每年约 2 - 3 倍增长，至 2027 年或超 10 亿美元。成本构成以加速器 / 服务器 / 互联折旧（47% - 67%）与研发薪酬（29% - 49%）为主，能源 2% - 6%。这迫使业界转向算法效率与底层架构的根本创新。

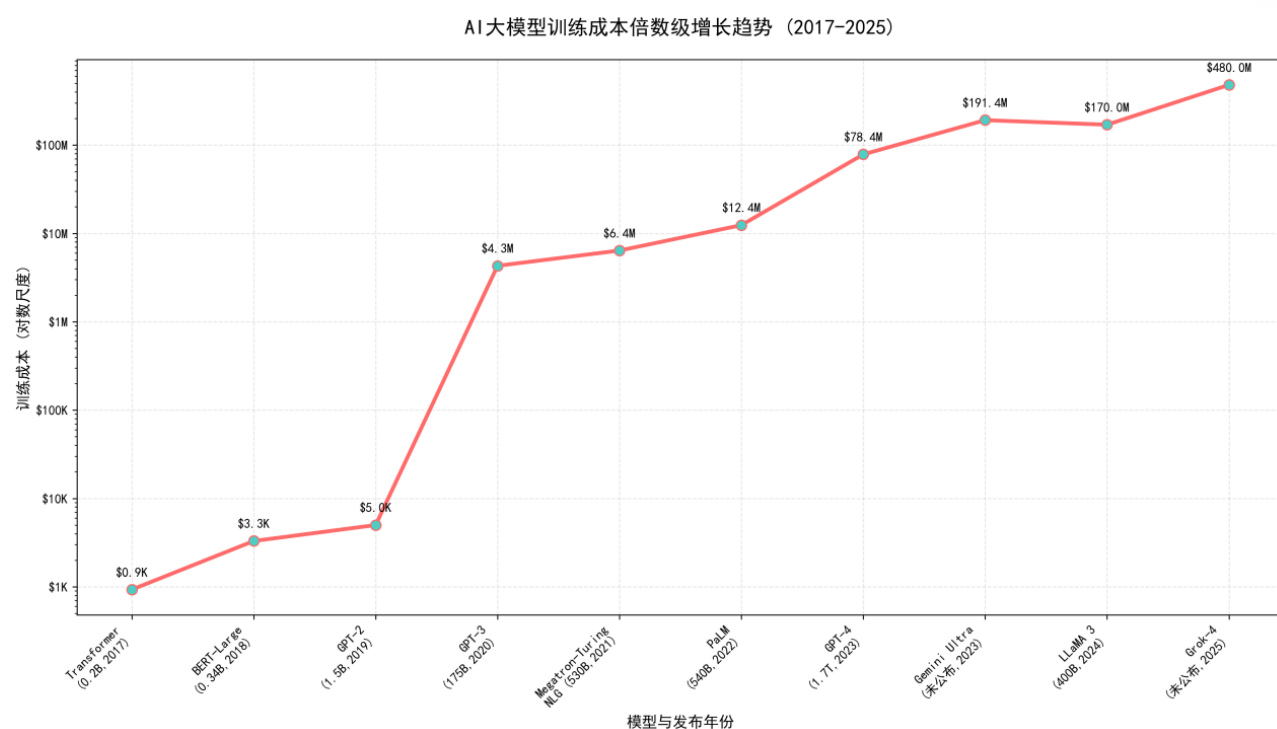


图 2.2 大模型训练成本增长趋势图（数据来源⁵⁶）

趋势四：AI 正迈向多模态与智能体的“复杂性”涌现

大模型技术正从单一模态向多模态融合，从简单的问答工具向具备复杂行为能力的智能体演进。Gartner 预测到 2030 年 80% 企业软件与应用将是多模态^[7]。多模态打通文本 / 图像 / 音频，智能体具备目标设定—推理—规划—工具调用能力，计算从可预测的“蛮力”转向动态、异构、有状态的“认知计算”，多模态打通文本 / 图像 / 音频，智能体具备目标设定—推理—规划—工具调用能力，计算从可预测的“蛮力”转向动态、异构、有状态的“认知计算”，基础设施重心转为低时延互联 + 智能编排平台。

2.3 大模型计算基础设施的挑战

模型发展的不同阶段，算力系统面临的瓶颈和挑战也在不断演变。AI 的核心任务从“生成”转向“交互”，即从一个数字世界的“内容创作者”，变为一个能够理解并影响物理世界的“行动者”。这一范式转变，对底层的计算架构提出了与单纯追求规模截然不同的新挑战。

1. 单卡阶段（CV 模型主导）：在计算机视觉（CV）模型为主的时期，模型可以完全放入单个加速卡中进行训练。此时的并行策略主要是数据并行，单卡的计算能力是主要瓶颈。

2. 单机阶段（小参数 NLP 模型主导）：随着 NLP 模型的出现，单卡显存不足以容纳整个模型，训练扩展到单机八卡。此时，数据并行和模型并行结合使用，节点内部的通信带宽（如 NVLink）成为瓶颈。

3. 传统服务器集群阶段（大模型主导）：当模型参数达到千亿乃至万亿级别，单机已无法满足需求，必须使用大规模服务器集群进行训练。这引入了序列并行、专家并行等更复杂的并行策略。训练集群的总规模（卡数）是数据并行（DP）、张量并行（TP）、流水并行（PP）和序列并行（CP）等多种并行维度的乘积。随着集群规模增大，单纯扩大数据并行（DP）维度会受到限制。因为在全局批次大小（GBS）固定的情况下，增加 DP 会减小每张卡的微批次大小（Num_BS），导致计算效率降低（产生 Bubble）。为了继续扩大集群，必须增加张量并行（TP）、流水并行（PP）和序列并行（CP）等维度（表 1）。然而，这些并行方式，特别是当并行域（如 TP>8）超出单台服务器的范围时，会产生大量且不可避免的跨节点网络通信。因此在这一阶段，跨服务器节点的通信带宽成为决定整体训练效率的核心瓶颈。

并行模式	主要通信方式	单次通信量	单步通信次数	能否被掩盖
数据并行（DP）	All-Reduce	单卡+GB以上	1次	是
张量并行（TP）	All-Gather、 Reduce- Scatter	GB级别	多次	否
流水并行（PP）	Send/Recv	MB级别	几十次	是
序列并行（CP）	All-Gather	十GB级别	1次	否
专家并行（EP）	All-To-All	GB级别	多次	否

表 1 不同并行模式下的通信特征

正是这种瓶颈的演变，最终凸显了传统服务器集群架构面临的三重系统性挑战。首先是通信墙，千亿级模型一次梯度同步即 TB 级数据，传统以太网难以承受。其次是功耗与散热墙，为破通信墙而提升密度，促使液冷、48V 供电成为标配。第三是复杂度墙：万级处理器带来故障常态化，从业界模型 GPT-3 (175B) 到 GPT-4 (1.9T) 的演进为例，随参数增至 10.8 倍，总集合通信达 34.1 倍，跨节点 RDMA 49.3 倍，光电转换 49.3 倍；任一组件或一次光电转换失败都会放大为全局可用度 / 利用率问题，进一步抬升运维复杂度。

2.4 小结

规模定律驱动的参数指数增长、从训练到推理的场景泛化、以及向多模态与智能体的跃迁，共同施压基础设施，形成通信墙—功耗散热墙—复杂度墙的循环。未来随着推理深入企业核心业务，需在低时延下实现高吞吐 / 高并发：逻辑推理等任务将带来推理算力需求或百倍增长，推理也将集群化，对互联与全局调度提出更高要求。因此，能否以更高效、更绿色的方式提供和使用算力，已不再是锦上添花的技术选项，而是决定 AI 规模化发展能否持续的战略挑战。

3.0 超节点的出现与演进

在大模型应用拉动下，传统数据中心的横向扩展范式暴露出跨机通信瓶颈。“一种以 AI 处理器高速互联为核心、实现跨节点大带宽 / 低时延的集成计算单元范式”初现端倪。尽管“超节点”目前是一个行业概念而非严格的技术标准，这标志着 AI 基础设施正在从“服务器集群”时代，迈向以软硬件一体化、绿色高效为核心特征的“集成计算单元”时代。

3.1 全球产业的演进路线：从硬件聚合到系统构建

过去的计算集群主要采用“横向扩展”（Scale-Out）架构，即通过通用以太网连接大量标准化服务器，这种模式能够很好地适配松耦合的计算负载。然而，大模型分布式训练对系统提出了截然不同的要求。其高频次的 All-Reduce、All-to-All 等集合通信操作，使得跨服务器的带宽与时延成为了根本瓶颈。随着集群规模的扩大，这一瓶颈愈发严重，严重制约了整体训练效率。

为突破传统“向外扩”（Scale-Out）架构的通信瓶颈，全球顶尖科技公司都在探索“向内聚”（Scale-Up）的超节点范式，但其实现路径各具特点。NVIDIA 的路径是以 NVLink 等专有硬件为核心，通过极致的垂直整合，将整机柜的 GPU 打造成一台逻辑上的“巨型单机”。而以华为昇腾为代表的国内企业，则通过“面向超节点的互联协议”创新，将大带宽低时延的互联范围从单机内部延伸至整个集群，实现跨主机的内存直接访问。2024 年 9 月的华为全联接大会上首次发布了昇腾 384 超节点以及超节点集群 Atlas 900 SuperCluster，实现了业界最大规模的高速总线互联。

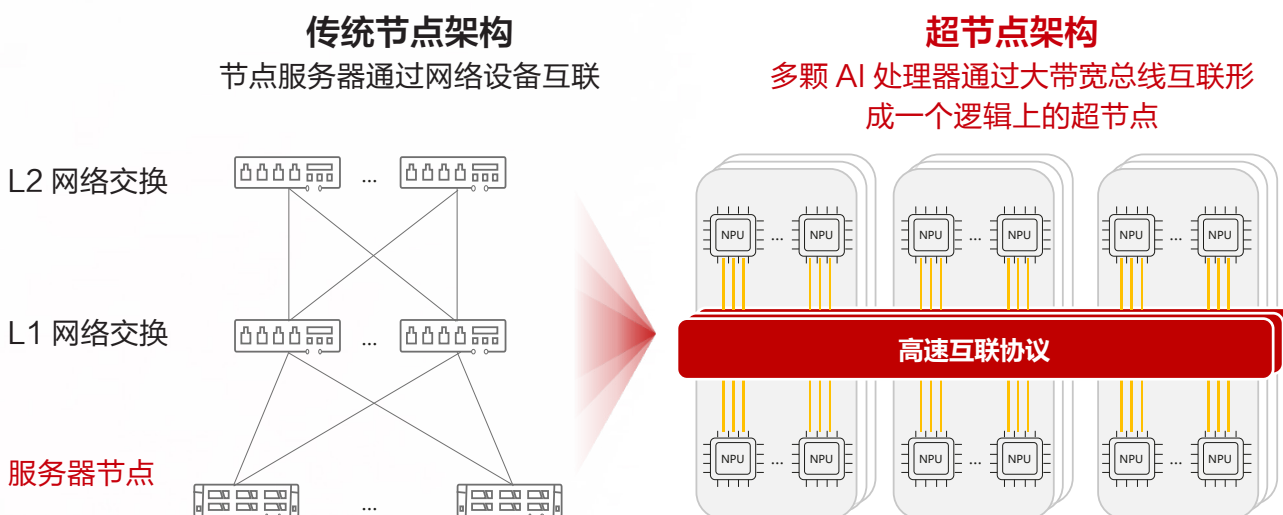


图 3.1 数据中心基础计算架构通信范式的演变

随着超节点概念的逐步成熟，其内涵也从最初聚焦于硬件互联，深化为软硬件一体化的全栈协同设计。一个现代化的超节点，其强大性能不仅来自先进的芯片和互联硬件，更依赖于一个被深度优化的、贯穿上下的软件栈，这个软件栈是硬件能力的“放大器”。这个体系包括多个关键层次，如硬件抽象与驱动层、专门优化的通信库、以及计算、内存等优化技术。这一从硬件到软件的全栈协同，是超节点区别于简单硬件堆砌的本质特征。

13.2 超节点技术产业生态发展格局

随着超节点成为 AI 基础设施的核心，围绕其构建的产业生态也呈现出两种截然不同但又相互竞争的发展模式。

第一种是“垂直整合”模式。其核心战略是通过对整个技术栈——从底层的 AI 处理器芯片，到核心的互联技术，再到上层的编程模型和软件库——端到端控制，来实现最优的性能和效率。在这种模式中，互联协议本身是闭源的，用户面临严重的厂商锁定风险，选择和灵活性受到极大限制，并且通常需要承担高昂的成本。

第二种是由其他芯片巨头、云服务商和大型科技公司组成的“协议开放”模式。其核心战略是以“兼容性”和“选择权”来对抗“封闭性”，旨在构建一个基于开放标准的、更加多元化的生态系统。这个联盟的主要参与者包括 AMD（推广其 ROCm 开源软件平台）、Intel（其 Gaudi 3 加速器内置标准以太网端口）以及 Meta、Microsoft 等超大规模用户。为了弥合开放生态的性能差距，多个旨在推动关键技术标准化的行业联盟应运而生，其中包括旨在增强以太网性能的“超级以太网联盟（UEC）”和旨在创建开放加速器互联标准的“超级加速器互联（UALink）”联盟。但此种模式下，缺少配套的软件实现，核心挑战在于生态系统的碎片化和性能优化的难度。

在中国，AI 基础设施的发展呈现出第三条路径——开源开放模式。以华为昇腾为代表的国内科技领军企业，正在加速构建从互联协议到基础软硬件全面开源开放的技术体系，建立生态系统。在面向超节点的互联协议开放的基础之上，通过一套软件架构满足不同代际超节点产品的持续演进，同时将核心计算架构 CANN、Mind 系列应用使能软件全面开源开放，结合对 PyTorch 等主流框架的全面兼容和体系化工具链，旨在最大限度地降低开发门槛，加速开发者和社区的融入，共建 AI 技术和产业生态。

13.3 小结

“超节点”的演进始于对传统 Scale-Out 架构通信瓶颈的突破，经历了从服务器到机柜级集成系统并最终深化为软硬件全栈协同的系统工程。它与 HPC、分布式计算等领域既有渊源又存在本质区别，其核心价值是为大模型深度优化的专用算力基石。超节点的形态和技术路线将持续演化，但其作为 AI 时代核心计算单元的地位已然确立。

4.0 超节点基础定义与特征

在人工智能大模型训练和推理等前沿技术的算力需求驱动下，传统分布式集群在通信效率、资源聚合能力上的局限性日益凸显，超节点凭借超高带宽互联、内存统一编址等技术特征，以及大规模灵活组网、高可靠运行等系统优势，成为支撑复杂计算任务的关键底座。

4.1 技术特征

技术特征方面，超节点的技术特征体系可划分为基础特征与扩展特征两大层次。其中，基础特征作为超节点产品的核心必备要素，构成产品运行的基石与核心竞争力；扩展特征则聚焦于产品能力的优化与延伸，涵盖可增强产品能力、及未来潜在扩展的技术特征，旨在通过持续技术创新提升超节点产品的功能完整性与市场适应性。

4.1.1 基础特征：大带宽、低时延、内存统一编址

超节点是 AI 计算节点通过高速互联协议组成更大内存空间的 AI 系统。超节点可以支持 32 及以上 AI 芯片，AI 芯片到交换芯片带宽不小于 400GB/s，交换设备时延小于 500ns。超节点域内 AI 芯片支持内存统一编址，AI 芯片使用内存语义可直接访问其他 AI 芯片的内存。

一、超大带宽和超低时延互联

超节点能够提供大带宽、低时延的互联能力。传统计算架构中，卡间互联依赖 PCIe 或以太网，跨服务器互联带宽多为 200~400Gb/s 且时延达数十微秒，在千亿参数模型训练的并行计算场景中，频繁的 GB 级数据通信阻塞，导致计算等待通信，成为性能瓶颈。

超节点借助高效的互联协议打破传统架构限制，支持更大规模 AI 处理器的高效协同，实现更大范围、更高流量的数据传输，从而突破系统性能。以昇腾 384 超节点为例，相较于传统服务器架构，通信带宽提升 15 倍、单跳通信时延从 2 微秒做到 200 纳秒，降低了 10 倍，在 DeepSeek、Qwen 等多模态、MoE 模型上，可以达到 3 倍以上的提升。

二、内存统一编址

超节点能够实现内存的全局管理和灵活访问。超节点内所有互联设备的内存地址需全局唯一，基于全局内存可实现任意设备间的灵活访问。这使得大模型训练中频繁的参数同步操作，无需经过传统的“序列化 - 网络传输 - 反序列化”流程，直接通过内存语义通信完成，提升小包数据传输及离散随机访存通信效率。

支持异步和同步两种模式访问。异步模式能够高效进行大块数据搬移（大块数据跨节点 DMA），适用于批量数据的快速传输与处理；同步模式则支持 Load/Store 内存语义访问 Cache Line 粒度小块数据，实现对小块数据的精细化处理需求。两种模式相辅相成，能够充分满足客户在不同应用场景下的需求。

●4.1.2 扩展特征：多级缓存池化、资源灵活配比

一、多级存储资源池化技术重构资源管理范式，突破系统间资源壁垒

超节点可通过资源池化技术将分散的计算、存储、网络资源抽象为统一逻辑资源池，以集中化管控的方式消除资源孤岛，实现动态弹性调度。

以大模型推理核心组件 KV Cache 为例，其访问效率直接决定推理延迟与并发处理能力。传统静态分配模式因资源隔离导致 KV Cache 重复存储与访问阻塞，难以应对高并发请求下的低延迟需求。资源池化技术通过构建“显存 - 内存 - 存储”多级存储体系，可将 KV Cache 从单机显存的物理限制中解放：单卡显存只需保留热数据，冷数据动态迁移至内存或存储层，从而支持百万 Token 级长上下文处理并实现批处理规模的倍级提升。

借助超节点大带宽和全局统一编址，通过支持不同 NPU 之间、不同 CPU 到 NPU 之间的 KV Cache 传输，缩短 KV Cache 传输路径和传输时间，压缩 KV Cache 的刷新时间，提升 KV Cache 命中率。为千 / 万亿参数模型在百万 Token 级长上下文与高并发请求场景下的稳定、低延迟推理提供保障。

二、资源灵活配比突破固定配比的刚性约束，实现资源精准按需配置

超节点可通过资源池化与软件定义架构的深度融合，将 CPU、NPU、内存、存储等物理资源解耦为可独立调度的资源池，根据任务特征自动调整各类型资源的配比比例：在计算密集型任务中提升 NPU 与内存占比，在访存密集型任务中提升带宽与显存资源占比，在存储密集型场景中增加 CPU 与存储资源配额。

例如，大模型训练场景中，需分配更多的 NPU 资源；推理场景中，则需要更高的带宽与显存资源；互联网推荐场景中，样本解析、特征处理任务对 CPU 资源需求较高。传统固定配比模式因资源绑定僵化，往往会导致资源闲置（如训练场景中存储资源过剩、推理场景中计算资源浪费），且无法应对任务负载动态变化。

资源灵活配比技术通过智能调度算法，可实现资源配比的动态调整，显著提升系统资源利用率以及任务处理效率，满足各类细分场景下差异化的资源需求。

4.2 系统特征

当前大模型技术的快速迭代，持续驱动智算基础设施向高密度、高协同方向升级。为构建更大规模的 Scale Up 组网，超节点在集群组网中需引入超节点通信域，依托定制化交换设备实现域内 NPU 高速互联（如图 4.1 所示），突破单机扩展的硬件限制，为大规模算力聚合提供架构支撑。此外，随着 Scale Up 组网规模的扩大与算力密度的提升，超节点的稳定性成为集群作业连续性的核心保障，需考虑器件、网络、系统等层面的可靠特性，以化解系统故障风险。在此基础上，针对单用户专属、多任务并行等差异化场景，超节点还需通过精细化资源调度、性能隔离与数据安全机制实现全场景适配，在满足复杂业务需求的同时最大化释放算力价值。

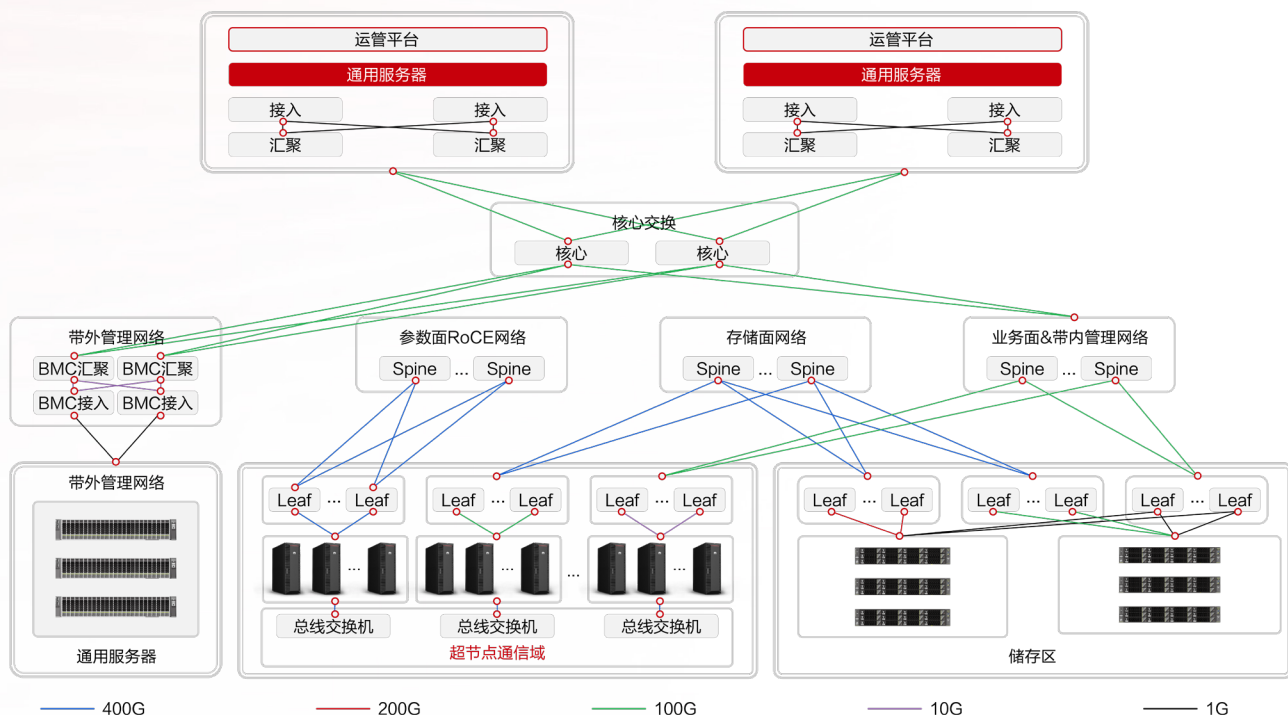


图 4.1 超节点集群组网架构（以昇腾 384 超节点为例）

为了更好地发挥超节点能力，超节点系统应具备以下特征：

4.2.1 超大规模

超节点作为新型算力基础设施，其大规模组网能力决定了模型训推效率与规模拓展的边界。以华为昇腾超节点产品为例（图 4.2），在 Scale Up 和 Scale Out 组网的技术协同下，超节点构建起高效、灵活的通信网络，成为突破算力瓶颈、支撑不同规模大模型训练及推理的核心技术路径。

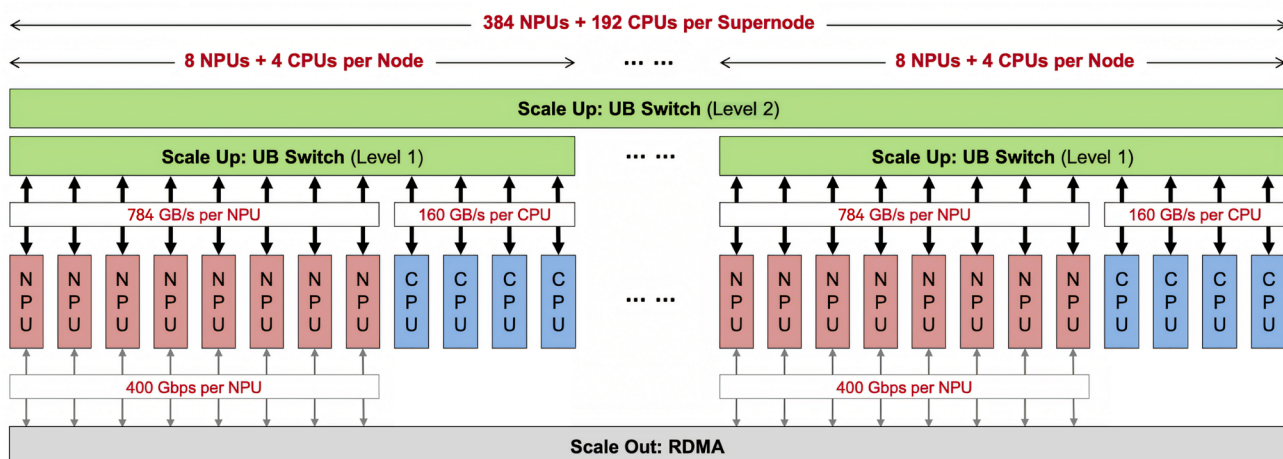


图 4.2 超超节点硬件架构示意图（以昇腾 384 超节点为例）

Scale Up 组网：突破单机算力边界，灵活构建大规模高速互联体系。

传统集群受限于 PCIe 带宽与拓扑结构，基于传统单机八卡架构实施张量并行（TP）、序列并行（CP）、专家并行（EP）等策略时，通信链路易达瓶颈（卡间带宽通常低于 100GB/s），且并行维度被限制在八卡以内。以大模型训练为例，当前主流的 MoE 模型（如 DeepSeek V3、Qwen3 等）都采用了 64 卡 EP 并行的方式，传统服务器形态，跨机通信存在瓶颈，优化困难，超节点的大带宽能够提升通信效率，缩短卡间不可掩盖的通信，降低模型性能调优难度，快速提升 MFU。

超节点需通过构建 Scale Up 高速互联体系实现规模破局：依托创新的互联技术打造大规模、高带宽、低时延的超节点域内通信环境，拓展并行域以适配各类并行场景，支持动态探索多维度并行策略组合（如灵活调整张量分片粒度与专家模型分配），挖掘任务最优解。同时，针对“大象流”与哈希不均导致的网络拥塞问题，超节点可通过包级负载均衡、乱序控制等技术动态平衡流量以提升域内带宽利用率，有效规避阻塞，充分释放通信性能潜力，为大规模并行计算提供高效算力聚合支撑。

伴随强化学习及 AI Agent 等训练技术的发展，训推一体的集群规模和节点间的流量持续提升，需要更大规模的超节点，消除卡间的通信瓶颈，并释放算力。

Scale Out 组网：实现集群化扩展，突破单节点算力限制。

面临万亿参数规模的超大型模型训练需求，超节点可通过叠加 Scale Out 组网，构建集群化组网架构。依托跨节点通信协议完成数据交互与任务调度，结合任务分配机制优化，将模型训练任务负载拆解至多个超节点并行执行，显著缩短训练周期，大幅提升模型训练效率，实现算力规模与训练效率的同步跃升。

●4.2.2 超高可靠

若要充分发挥超节点算力，稳定性是决定系统计算效率及成本的重要指标之一，最大程度保障训练任务不中断，训练数据和结果不丢失。超节点稳定而可靠的运行依赖可靠的器件、可靠的网络以及可靠的系统。

可靠器件构筑物理层稳定基石：

超节点硬件器件的可靠性是系统稳定运行的核心前提。在大模型训练过程中，集群高负载运行带来了持续的高功耗，极易引发核心器件温度剧烈波动，导致故障发生概率激增。据《The Llama 3 Herd of Models》研究显示，Llama3.1 模型在万卡集群训练时，平均 3 小时便会出现一次故障，其中绝大多数故障源于器件层面的失效。

为应对这一挑战，超节点系统需要深度考虑硬件可靠性，覆盖器件的全生命周期(如图 4.3 所示)：生产阶段通过高低温、震动、冲击等极限压测方法筛选合格产品；选型阶段选用超宽温度、稳定性能、低失效的电信级器件；使用阶段部署液冷散热技术，精确控制核心器件工作温度，降低因热应力导致的故障风险。这些措施可使超节点系统在满负载运行时的器件级故障率大大降低，为千亿参数模型连续训练提供硬件级可靠性保障。

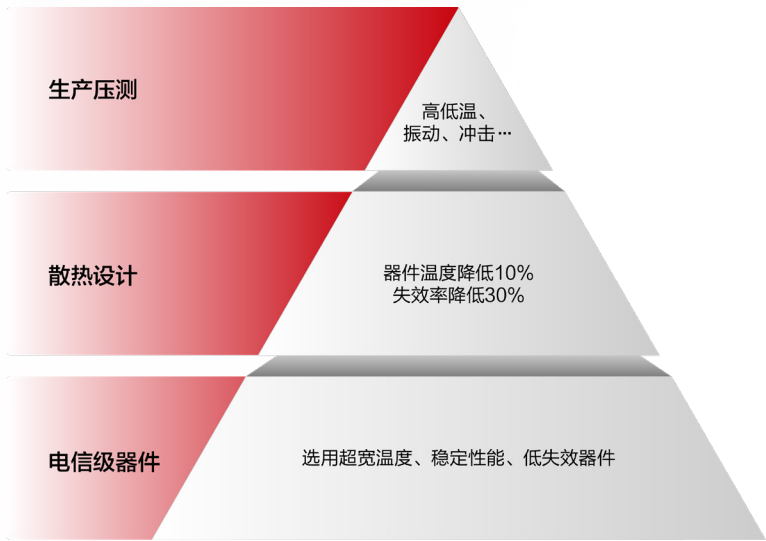


图 4.3 硬件可靠性策略

可靠网络保障数据传输无中断：

超节点网络的可靠性是支撑大规模计算任务的核心要素。如图 4.4 所示，超节点网络涵盖超节点通信域与跨节点通信域，相比传统服务器系统，超节点集群光模块的数量增加了约 7 倍以上，面临更高的网络失效率的挑战。

在冗余节点备份方面，由于超节点通信域内带宽显著高于跨域带宽，当域内计算节点故障时，需设计域内冗余节点配置与整节点迁移机制，避免因带宽不匹配导致故障无法及时恢复。针对光模块数量庞大且张量并行（TP）域高度依赖其进行大流量数据传输的问题，超节点系统需具备光模块故障检测与动态修复机制，当光模块发生局部故障时自动切换至冗余通道，保障数据传输链路的连续性。针对多平面通信架构，超节点应设计链路收敛算法，在某一平面链路故障时快速路进行路径切换与流量重定向，大幅缩短故障修复时间。

这些机制共同保障了超节点在大规模计算场景下的稳定运行，为高效计算提供了可靠支撑。

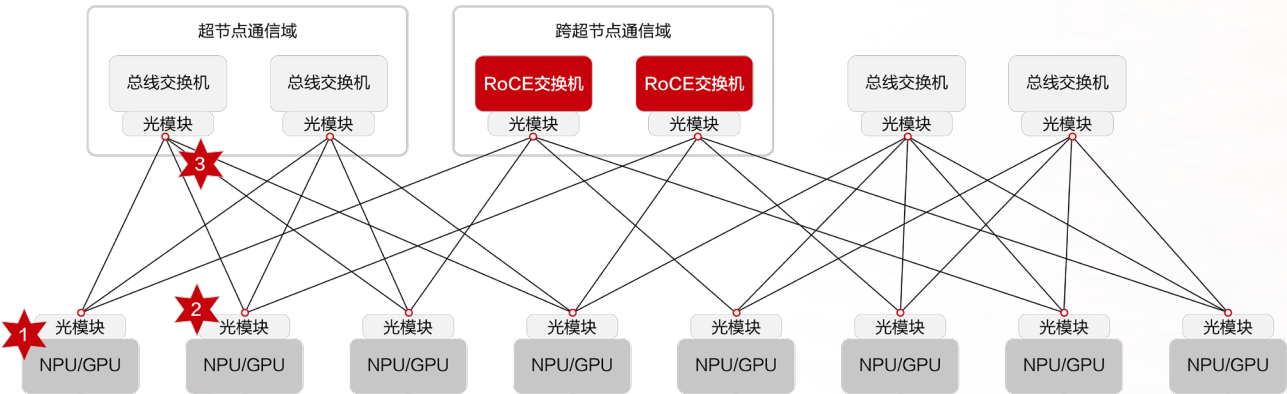


图 4.4 超节点通信域(以昇腾 384 超节点为例)

可靠系统实现全域统管与故障排除：

随着超节点系统规模扩张与架构升级，传统多系统独立运维、协议碎片化的管理模式已难以满足需求，尤其超节点通信域及新型交换设备的引入，进一步加剧了系统运维管理与故障排查的复杂性。故障主动预防与恢复能力是超节点系统运维的关键能力。

一、故障主动预防，集成多种故障机理模型与预测算法，构建智能化故障预防体系。通过对器件运行数据的挖掘与趋势分析，系统能够提前感知器件的亚健康状态，预判潜在故障风险。针对可能出现的性能下降、连接异常等问题，系统可主动触发预防性维护措施，有效避免因核心器件失效导致的系统性能瓶颈与业务中断，确保集群系统长时间稳定运行。

二、分级故障恢复，基于故障影响范围与业务特性，制定阶梯式恢复策略。对于算子级故障等影响范围较小的故障，可优先采用在线恢复策略，通过快速重启业务实现秒级恢复；对于节点故障、作业故障等影响范围较广的故障，系统可优先调用超节点内备份节点接管任务，故障修复后的节点则重新纳入备份资源池；当超节点内部缺乏可用备份节点时，系统将启动逻辑超节点迁移机制，将故障节点业务快速迁移至其他可用节点，实现分钟级恢复。这种分级恢复策略确保不同类型业务在故障发生时均能被实时高效地处理，以最小时间代价保障业务连续性。

●4.2.3 灵活切分

超节点实现大规模组网和稳定运行的同时，所面临的场景也从单任务独占到多任务并行逐步多元化，不同场景对资源调度、性能隔离、数据安全要求差异显著。因此，构建适配全场景的系统能力，成为超节点释放算力价值、支撑企业数字化转型的必然选择。

单任务场景下的资源独占与性能优化：

针对单用户单任务场景（如千亿参数大模型训练、高精度科学计算）对算力连续性和资源纯度的高要求，系统应将全部计算、存储和网络资源集中调配至单一任务，避免资源争用与调度损耗。用户仅需关注超节点域的大小和模型切分策略，利用超节点互联优势即可实现计算性能的最大化释放。

多任务场景下的逻辑切分与高效协同：

面对多任务并行（如同时运行模型训练、在线推理、数据预处理）的混合负载场景，超节点应采用逻辑切分技术，将物理超节点拆分为多个逻辑超节点，各逻辑节点间通过高速互联协议实现低时延通信（如图 4.5 所示），允许系统根据不同任务的算力需求、优先级和资源偏好进行差异化调度。同时，逻辑切分技术赋予系统精准定位和隔离故障的能力，当某个逻辑超节点出现异常时，仅影响该节点承载的任务，其余逻辑节点仍可保持正常运行，大幅降低故障扩散风险。



图 4.5 逻辑超节点虚拟切分示例

5.0 超节点应用案例

5.1 支撑大模型创新及云服务场景

面向大模型训练、迭代调优及推理云服务能力场景需求，超节点产品可充分发挥高带宽低时延、强协同、高效率等特点，作为最佳亲和 MoE 模型架构的系统，成为大模型创新的重要基础设施，同时也为大模型工程优化带来了更多的可能性，可以为互联网、大模型厂商等需要前沿模型探索及深度工程创新的客户带来更好的支撑。

互联网基础模型的探索向着更多参数(万亿级别)、更长序列(512k->1M)、更多专家(200->400)的方向发展，需要 TP、CP、EP 等多维并行，对计算平台的集群通信能力有很大挑战，昇腾 384 超节点的系统创新匹配了模型发展的趋势，互联网行业正基于 384 超节点在 LLM 推理、预训练、强化学习、多模态、生成式推荐等多个方向进行技术创新。与此同时，互联网基于大模型的 AI 应用也处于高速增长阶段，希望有更优的推理成本和更好的用户体验，通过大规模专家并行推理方案，结合 384 超节点的系统通信优势，以 DeepSeek R1 为例，单卡吞吐可以达到 2000tokens/s，时延低至 15ms，可以很好地满足互联网推理要求。

互联网客户 1 基于超节点部署大规模专家分布式并行的推理方案。在 MoE 大模型推理场景中，实现 40~50% 单 token 成本降低，相比业界，单卡吞吐提升 2.4~2.8 倍，业务时延从 30ms~50ms 降低到 15ms，极大程度提升用户体验、降低推理成本。

互联网客户 2 基于超节点的集群通信能力进行 MoE 基础模型预训练。对多专家和长序列带来更好的支持，整个训练时间中未掩盖通信的时间占比从业界的 35% 优化至 15%，大幅提升 MFU，支持客户训练出业界领先的基础模型。

互联网客户 3 基于超节点开展强化学习训练。结合超节点的性能和对强化学习主流技术生态的支持（如 verl+vLLM）快速支持主流模型（如 Qwen、DeepSeek）和算法（GRPO、DAPO 等），支持多个客户快速探索强化学习算法创新。

互联网客户 4 基于超节点加速多模态模型创新。探索 DiT-MoE、生成理解统一、世界模型等前沿的技术方向，已支持客户在生成和理解领域持续孵化出 SOTA 模型并在业务中上线商用。

运营商基于超节点部署企业训练及推理智算服务业务，支持企业前沿大模型预训练、行业 / 场景大模型快速迭代、提供推理 API 服务产品等。2025 年 4 月 26 日，中国电信全球首个商用智算昇腾超节点在粤港澳大湾区（韶关）算力集群正式发布上线，并依托自研智算测试与适配优化平台，针对基准性能、典型场景及资源配置，基于昇腾 384 超节点开展推理性能自动化测试与系统级调优，

使能中国电信根据业务场景灵活配置资源及并发策略，最大限度发挥系统性能，有效服务于复杂推理、多智能体协同等新场景；同时，中国电信基于自研星辰大模型、DeepSeek 等业界前沿基础大模型，逐步探索预训练、强化学习、微调等技术创新。目前，智算昇腾超节点服务产品已基于中国电信“息壤”一体化服务平台在天翼云官网上线，面向互联网企业、大模型厂商、科研 / 教育 / 医疗创新机构、央国企等提供低延时、高吞吐的算力服务。

5.2 加速人工智能科学计算，服务算法创新

超节点提升超算中心人工智能计算场景能力，加速 AI 赋能科学研究。通过引入超节点先进技术，赋能科学研发场景，如在大气海洋环境领域，对大气、海洋等各种活动进行建模和分析，预测气候变化，防范和减轻气候变化带来的破坏；在生物医药健康领域，帮助用户分析和处理生命科学中的海量数据，加快我国生命科学的研究进展；在工业设计领域，使工程师能在几分钟或几小时内仿真和测试数千种设计方案，为我国科技发展提供发展动力。

目前超节点技术已落地 AI4S 场景，提升科学场景研发效率。国家超级计算广州中心国内首批部署昇腾超节点，满足不断增长的智算需求。该超节点采用昇腾 AI 基础软硬件，结合高速互联协议，具备大规模、高性能、高可用等特性。在推理场景中，通过大规模专家并行推理等技术，单卡的吞吐量提升了 20 倍，为科研人员和企业提供更高效、更强大的计算服务。

超节点支撑大模型技术研究，助力我国人工智能模型算法创新。鹏城实验室基于昇腾超节点开展大模型技术研究，通过昇腾软硬协同的创新架构实现了科研效率的全面提升。实验室构建了国内首个 EFLOPS 级 AI 算力集群“鹏城云脑”，在昇腾超节点的支持下，“鹏城云脑”的训练效率将大幅提升，原本 MoE 模型训练需要 3 周时间，借助昇腾 384 超节点，时间将压缩至 1 周；基于昇思 MindSpore 进行了高性能的监督微调和 GRPO 强化学习训练框架的构建，发布开源数学推理模型 PCL-Reasoner-V1，率先突破大模型强化学习技术全栈壁垒，相关研究成果在国际顶级学术会议中发表，展现我国人工智能算法创新能力。同时，昇腾 384 超节点采用液冷技术，数据中心 PUE 低至 1.1，鹏城云脑实现成本运营降低，实现绿色高效运营，较低的能耗符合绿色数据中心的发展理念，树立良好的社会形象。

5.3 助力行业企业智能化升级

与行业集团大型业务场景需求融合，为全行业应用场景智能化需求提供坚实底座。人工智能是促进企业跨越式发展、加速产业优化升级、带动生产力整体跃升的核心战略资源。当前传统行业智能

化转型面临业务场景需求激增、低资源算力重复建设等问题，超节点提供高算力、大带宽、低时延能力，能够满足行业大型集团多业务场景、高并发等人工智能场景推理需求。

能源行业率先布局超节点基础设施，服务电网生产、企业经营、客户服务、设备管理等业务领域。某电力行业企业作为能源行业龙头企业，积极打造人工智能基础设施，目前已建成国产自主可控的两级智算中心，算力总规模达 600+PFLOPS，基于先进的算力基础设施，打造了全尺寸、多模态、全自主的光明电力大模型，是首个通过中国信通院、电子标准院国家权威机构双认证的行业大模型，也是首个通过国资央企大模型测评体系认证的电力行业大模型，专业能力达到卓越水平，较基座模型平均提升 15%。为进一步提升技术先进性、服务集团业务发展需求。企业计划于 2025 年在行业内率先完成超节点集群建设，形成一体化算力体系，实现全网算力资源统筹纳管和集约使用。在智能客户服务、供电方案智能生成、电网智能诊断分析、智慧办公、主网负荷智能预测等场景中，基于昇腾大规模专家并行技术系统级优化，实现超高的系统吞吐性能，满足各业务部门的高并发 AI 服务调用需求。

在制造业场景，行业头部客户已规模部署昇腾 A3 超节点，通过深度优化，超节点在模型训练及微调中的性能表现比传统 AI 服务器产品提升 2.5–3.5 倍，客户通过采用昇腾超节点的分布式并行切分策略，算力利用率（MFU）也由 37.3% 提升至 55.9%。后续超节点将作为企业智能化升级的强劲算力引擎，在大模型训练、AI 智能体、代码辅助、AI 问数、工厂大脑、供应链管理等场景中持续支撑客户取得创新突破。

在交通行业，国内某港口近年来通过数字化、智能化技术深度应用，加速向智慧绿色港口转型。为了进一步加速智能化战略进程，建设了全球首个基于昇腾超节点算力底座的数智化港口行业应用基地。梳理 6 大领域 18 类共 80+ 场景，重点对准智慧码头生产、平安港口、智慧管理 3 大核心领域展开创新，首批试点落地船舶智能排泊、干散货皮带智能巡检、码头设备智能运维等智能化应用场景。

在运营商，中国联通在呼和浩特云数据中心落地昇腾 384 超节点，结合“联通星罗”先进算力调度平台，为人工智能高质量发展注入新动能。面向内部自用，部署元景、DeepSeek 等主流大模型，支撑智慧办公、市场营销、运营维护、网络安全等场景。此外，超节点可服务于政教医工等领域，提供大模型极致训推服务，使能行业用户发展。浙江移动引入昇腾 384 超节点先进算力，通过超节点大规模、高带宽、低时延、算力切分、虚拟化及训推一体等一系列领先能力，结合 MaaS 平台，面向企业内部提供统一的大模型 API 服务，匹配不同应用场景对训、推算力的需求，服务于智慧运营、智慧管理及智慧运维等核心场景，支撑 BOM 核心业务支撑系统的多业务领域的智能化转型。

本报告系统性地论证了，在生成式 AI 浪潮的驱动下，传统数据中心架构正面临通信、功耗和复杂度这三重不可逾越的壁垒。作为应对这一系统性危机的革命性方案，AI 超节点通过高速互联技术架构创新，并辅以软硬件全栈协同的深度优化，正在重新定义 AI 时代的算力基石。围绕超节点的产业格局已经形成，AI 超节点的形态和技术路线将持续演化，其作为 AI 时代核心计算单元的地位已然确立。面向未来，大模型技术会加速迭代，围绕超节点的软硬件技术也将持续演进，带来如下几个方面的变化：

一、超节点网络时延向纳秒级演进

超节点的低时延从“微秒级”到“纳秒级”，比当前主流网络时延降低一个量级；大带宽实现任意节点间突破 Tbps 级，构筑数据中心高速公路；通过多样化集群拓扑优化，实现超节点从“机柜级”到“园区级”乃至跨 DC 大模型训推；更低成本和多样化的组网方案会加速演进，并逐步成为行业标准。

二、算力密度推动液冷全面落地

单柜算力密度数量级提升，迈入“兆瓦级整机柜”时代，热 - 电 - 结构三位一体实现液冷超节点将从“可选项”变为“交付默认配置”。

三、负载解耦与异构编成为超节点刚需

AI 负载及部署范式快速变化，PD 分离、AFD 分离、多模态 VAT 和 LLM 分离等范式对 AI 算力芯片有着不同要求，超节点通过灵活异构算力资源编排达成综合性能最优化部署。

四、多样性算力池化

超节点会带动多样性算力池化资源协同计算，通过超节点互联协议为依托支持各类 XPU 池化，实现对等访问，系统资源高效协同计算。

五、基于超节点的原生创新模型加速涌现

基于超节点下的模型创新迭代加速，TOP 级开源模型将“为超节点而生”，亲和性设计原生支持“机柜 - 模组 - 芯片”多级并行并支持自动最优并行维度搜索，实现开发者零感知。

六、极致低时延引爆应用

超节点成为下一代“爆款 APP”的默认底座，交互式 Agent 如 Sora-class 视频生成、实时 3D AIGC、AI NPC 对战等场景支持用户体验从“可用”到“沉浸”式转变；超节点成为短视频、直播电商、元宇宙社交等“杀手级应用”的准入门槛。

除此之外，超节点架构的持续发展离不开稳健的生态支持，基础协议及软硬件的开源开放，是加速超节点技术及产业生态发展的必然方向。未来，随着超节点向“光算存网”一体化演进，基础协议与软硬件的开源开放还将进一步深化：例如光互连协议的开源将实现“电－光协同调度”技术共享，持久内存与内存池的开源适配将推动超节点存储成本持续下降。这种“共建共享”的生态模式，最终将让超节点技术从“头部企业专属”走向“全产业可用”。

总而言之，AI超节点的演进之路远未结束。它将不断吸收和融合最前沿的计算、网络 and 物理技术，持续突破性能、效率和规模的边界，为人类迈向通用人工智能的伟大征程，提供坚不可摧的算力基石。

1. Noffsinger, J., Patel, M., Sachdeva, P., Bhan, A., Chang, H., & Goodpaster, M. (2025, April 28). The cost of compute: A \$7 trillion race to scale data centers. McKinsey & Company. <https://www.mckinsey.com/industries/technology-media-and-telecommunications/our-insights/the-cost-of-compute-a-7-trillion-dollar-race-to-scale-data-centers>
2. Kaplan, J., McCandlish, S., Henighan, T., Brown, T. B., Chess, B., Child, R., ... & Amodai, D. (2020). Scaling laws for neural language models. arXiv. <https://arxiv.org/abs/2001.08361>
3. 我国已有 433 款大模型完成备案并上线提供服务——AI 产业迈入规模化应用新阶段 - 中国网, http://www.china.com.cn/txt/2025-07/01/content_117956489.shtml
4. Cottier, B., Rahman, R., Fattorini, L., Maslej, N., & Owen, D. (2024). How much does it cost to train frontier AI models? arXiv:<https://doi.org/10.48550/arXiv.2405.21015>
5. Visualizing the Training Costs of AI Models Over Time. Visual Capitalist. <https://www.visualcapitalist.com/training-costs-of-ai-models-over-time/>
6. Artificial Intelligence Index. (2025). Artificial intelligence index report 2025. Stanford University.
7. Gartner. (2025, July). Gen AI data snapshot. Gartner. <https://www.gartner.com/cn/newsroom/press-releases/july-2025-gen-ai-data-snapshot>

■ 1. 核心概念

AI 处理器 (AI Processor) :

本文中提及的 AI 处理器泛指用于人工智能计算的加速器，如神经网络处理器 (NPU) 和图形处理器 (GPU) 等。

规模定律 (Scaling Law):

揭示了 AI 模型性能与参数量、数据量、计算投入之间存在幂律关系的法则。该法则驱动了模型参数的指数级增长，是推动基础设施走向超节点架构的核心理论之一。

超节点通信域 (Supernode Communication Domain):

在超节点集群组网中，为实现大规模 Scale Up 组网而引入的特定通信区域。在该域内，通过定制化的交换设备实现 NPU 之间的高速互联。

逻辑超节点 (Logical Supernode):

通过逻辑切分技术，将一个物理超节点拆分为多个逻辑上隔离的计算单元。每个逻辑节点可以承载不同的任务，从而实现在多任务并行场景下的高效协同、性能隔离和故障隔离。

■ 2. 技术架构与特征

内存统一编址 (Unified Memory Addressing):

超节点的一项核心技术特征，指所有互联设备的内存地址全局唯一，可基于全局内存实现任意设备间的灵活访问。这使得跨设备的数据同步能通过内存语义直接完成，无需传统的网络传输流程，从而提升通信效率。

资源池化 (Resource Pooling):

将分散的计算、存储、网络等物理资源抽象为统一的逻辑资源池的技术。它通过集中化管控和动态弹性调度，消除资源孤岛，提升资源利用率。

资源灵活配比 (Flexible Resource Allocation):

超节点通过资源池化和软件定义架构，将 CPU、NPU、内存等物理资源解耦，并根据任务特征自动调整各类资源的配比。例如，为计算密集型任务增加 NPU 占比，为推理场景增加网络与显存资源。

■ 3. 互联与通信

灵衢总线 (Lingqu Bus):

华为提出的一种通用互联协议，以网络为中心、消息驱动，通过分配全网唯一的逻辑地址来支持跨主机访问，并可封装于 IP/UDP 协议以适应大规模池化需求。

Scale Up:

指在单一系统内通过高速互联扩展 AI 处理器规模的方式。在超节点中，主要指通过定制化互联协议构建大规模、高带宽、低时延的内部通信环境，以突破单机算力边界。

Scale Out:

指通过网络连接多个独立节点来扩展算力的方式。在超节点架构中，通常指将多个超节点单元组合成一个更大规模的集群，以满足万亿级超大模型的训练需求。

■ 4. 系统与运维特征

RAS (Reliability, Availability, and Serviceability):

即可靠性、可用性与可服务性体系。在超节点中，RAS 是应对万级处理器规模下故障常态化的关键系统能力，涵盖了从可靠器件、可靠网络到可靠系统的全方位设计。

分级故障恢复 (Hierarchical Fault Recovery):

基于故障影响范围和业务特性制定的阶梯式恢复策略。对于算子级等小故障可实现在线秒级恢复；对于节点或作业级故障，则通过调用备份节点或逻辑迁移，实现分钟级恢复。

近存计算 (Near-Memory Computing):

将部分计算逻辑直接集成到存储芯片（如 HBM）内部或其附近的设计思想，旨在解决 AI 工作负载面临的“内存墙”瓶颈。

I 5. 常见缩略语

AI: 人工智能 (Artificial Intelligence)

CPU: 中央处理器 (Central Processing Unit)

GPU: 图形处理器 (Graphics Processing Unit)

NPU: 神经网络处理器 (Neural Processing Unit)

EP: 专家并行 (Expert Parallelism)

HBM: 高带宽内存 (High Bandwidth Memory)

KV Cache: 在大模型推理中用于存储上下文信息的关键 - 值缓存

LLM: 大语言模型 (Large Language Model)

MaaS: 模型即服务 (Model as a Service)

MoE: 混合专家模型 (Mixture of Experts)

PCIe: 一种高速串行计算机扩展协议标准

PUE: 电源使用效率 (Power Usage Effectiveness)，衡量数据中心能源效率的指标

RDMA: 远程直接内存访问 (Remote Direct Memory Access)

ROI: 投资回报率 (Return on Investment)

TCO: 总体拥有成本 (Total Cost of Ownership)

TP: 张量并行 (Tensor Parallelism)

商标声明

 HUAWEI, HUAWEI,  是华为技术有限公司商标或者注册商标，在本手册中以及本手册描述的产品中，出现的其它商标，产品名称，服务名称以及公司名称，由其各自的所有人拥有。

免责声明

本文档可能含有预测信息，包括但不限于有关未来的财务、运营、产品系列、新技术等信息。由于实践中存在很多不确定因素，可能导致实际结果与预测信息有很大的差别。因此，本文档信息仅供参考，不构成任何要约或承诺，华为不对您在本文档基础上做出的任何行为承担责任。华为可能不经通知修改上述信息，恕不另行通知。

版权所有 © 华为技术有限公司 2025。保留一切权利。
非经华为技术有限公司书面同意，任何单位和个人不得擅自摘抄、复制本手册内容的部分或全部，并不得以任何形式传播。