



专题：智算网络

智算互联综述

张云勇¹, 闫硕¹, 陈永铭², 张启明²

(1. 中国联合网络通信有限公司云南省分公司, 云南 昆明 650206;

2. 中兴通讯股份有限公司, 云南 昆明 650034)

摘要: 随着大模型参数量突破万亿规模, 智算互联面临超大规模组网、低时延通信、高带宽同步等技术挑战。研究构建了包含吞吐量、时延、扩展比等指标的多维评价体系, 分析了大模型训练、人工智能 (artificial intelligence, AI) 推理和边缘计算三大应用场景的需求特点。通过对比主流科技企业的解决方案, 总结了 CLOS 架构、Fat-Tree 拓扑等创新实践, 重点探讨了互联协议、网络拓扑、拥塞控制等关键技术, 并展望了开放协议、光电融合等未来发展方向。研究表明, 智算互联技术的持续创新将为 AI 发展提供关键基础设施支撑。

关键词: 人工智能; 智算互联; 大模型训练; 网络拓扑; 光电融合

中图分类号: TN92

文献标志码: A

doi: 10.11959/j.issn.1000-0801.2025165

A survey on intelligent computing interconnection

ZHANG Yunyong¹, YAN Shuo¹, CHEN Yongming², ZHANG Qiming²

1. Yunnan Branch of China United Network Communications Co., Ltd., Kunming 650206, China

2. Zhongxing Telecommunication Equipment Corporation, Kunming 650034, China

Abstract: As model parameters surpass the trillion-scale mark, intelligent computing interconnection faces technical challenges including ultra-large-scale networking, low-latency communication, and high-bandwidth synchronization multidimensional evaluation. The framework incorporating key metrics such as throughput, latency, and scaling ratio was established, the distinctive requirements of three major application scenarios: large-scale model training, AI inference, and edge computing, were analyzed. Through comparative analysis of solutions from leading technology enterprises, innovative practices were summarized including CLOS architecture and Fat-Tree topology, with discussions focused on critical technologies like interconnection protocols, network topologies, and congestion control. Future development directions such as open protocols and optoelectronic integration were also outlined. The findings demonstrate that continuous innovation in intelligent computing interconnection technologies will provide crucial infrastructure support for AI development.

Key words: AI, intelligent computing interconnection, large-scale model training, network topology, optoelectronic integration

收稿日期: 2025-04-28; 修回日期: 2025-05-25

入选第二十七届中国科协年会“AI时代网络技术创新”专题

0 引言

人工智能 (artificial intelligence, AI) 技术的范式革新正以前所未有的速度重塑计算生态。从 Transformer 架构突破语言建模瓶颈^[1], 到以 GPT-4^[2]、DeepSeek^[3]为代表的大语言模型 (large language model, LLM) 实现跨模态推理, 模型参数量已突破万亿规模, 训练算力需求呈指数级增长。研究表明, 单次 LLM 训练需消耗数千万图形处理单元 (graphics processing unit, GPU), 成本达到千万美元量级^[4], 这催生了由异构加速器 (GPU/TPU^[5-7]) 构成的超大规模算力集群, 并推动分布式计算成为 AI 基础设施的核心形态。

智算互联技术通过构建集群内部的计算拓扑与通信协议, 实现分布式异构算力资源的聚合。在硬件层面, 优化芯片间互连路径, 使单集群内 GPU 间通信时延降低至纳秒级; 在系统层面, 支持多任务并行处理。然而, 随着模型规模的不断增长, 算力需求也不断提升, 智算互联的重要性日益凸显。

(1) 智算互联的挑战

智算互联目前面临诸多挑战, 主要包括以下几方面。

① 随着模型参数的增加, 对算力和网络规模的需求都在急剧增加, xAI 公司已经搭建了超 10 万 GPU 的算力集群, 未来网络规模可能更大, 超大规模的智算互联如何高效运行正在成为工业界关注的技术难点。

② 在并行计算中, 计算与通信很难完全掩盖, 通信时延决定了并行计算的规模, 因此智算互联对于时延的要求越来越高, 传统的互联方案面临巨大的挑战。

③ 智算互联的单端口带宽及规模都在增加, 对于单交换芯片的交换容量提出了更高的要求, 大容量交换芯片可以简化网络结构, 但是大容量交换芯片对于芯片制造工艺要求更高, 大容量交

换带来了高功耗, 散热也存在比较大的挑战。

④ 大模型训练中的故障率较高, Meta 在训练 Llama 3 过程中, 使用 NVIDIA H100 GPU 集群 (16 384 卡) 在全天无间断训练时间期间 (54 天) 内出现了 419 次意外故障, 平均每 3 h 就有 1 次^[8]。因此超大规模智算互联的可靠性至关重要。

(2) 智算互联的核心评价体系与指标

智算互联对于网络传输的带宽、时延等提出严苛要求, 因此智算互联的核心评价体系需围绕以下维度及指标构建。

① 性能维度: 算力协同的实时性保障

吞吐量: 指单位时间内网络传输有效数据的最大值, 直接决定互联的性能。在大模型训练中, 单次梯度同步需传输 TB 级参数 (如 GPT-4^[2] 单次迭代同步 1.8 万亿参数, 数据量超 10 TB)。若网络吞吐量不足, 通信时间占比将激增至 50% 以上, 导致 GPU 计算单元频繁空闲^[9]。

端到端时延: 指数据从发送端到接收端的全程传输耗时, 时延决定多卡并行协作的实时性。时延的增加对于同步的时效性会产生明显影响^[10]。

② 效率维度: 资源利用的最优化

带宽利用率: 反映实际传输数据量与理论带宽上限的比值, 衡量网络承载效率。出色的负载均衡以及拥塞控制算法在相同场景下可以显著提升带宽利用率^[11]。

③ 扩展维度: 规模增长的可持续性

扩展比: 指集群规模扩大 N 倍时任务执行性能提升的比率, 表征算力扩展效率。理想状态下, 期望集群实现接近线性的扩展比, 但实际中各种因素导致的性能瓶颈会极大地降低扩展比。智算互联通过各种优化手段消除互联环节瓶颈, 为提升加速比创造良好的条件^[9]。

1 典型工作场景与主流厂商解决方案

在智算互联的生态体系中, 典型工作场景可



依据计算需求与资源分布特征划分为三大类：大模型训练场景、AI推理服务场景与边缘AI场景，三者共同构成分层协同的智能化算力网络。它们通过智算互联平台形成“集中训练-分布式推理-边缘协同”的闭环，依托统一资源编排、跨域数据流管理与智能调度算法，实现算力资源的最优配置与全场景AI服务的高效供给。

1.1 大模型训练场景

AI大模型的扩展定律^[12]和涌现能力^[13]驱动AI大模型规模的持续增大。当前以Transformer及其改进架构为代表的大模型参数规模增加至万亿级别，其训练任务所需的算力和存储资源已经远超单个智算服务器的能力。这需要通过高性能智算互联构建包含千卡、万卡的高性能分布式计算集群，并在各种并行工作策略的调度下协同完成海量语料数据的训练任务。该场景下在算力扩展的同时需解决跨节点通信带宽瓶颈与训练稳定性问题，在训练过程中，需要动态调度异构算力资源，结合梯度压缩、稀疏训练等技术降低能耗与成本，并通过容错机制应对硬件故障导致的训练中断风险。

1.2 AI推理服务场景

随着以DeepSeek R1^[14]、OpenAI O1^[15]、阿里巴巴QWEN^[16]为代表的基础大模型综合能力持续提升，大模型推理应用也逐步从实验室走向商业化。大语言模型推理相对于训练场景，更加关注智算互联的成本和响应速度。应用于推理场景的智算互联技术需具备如下技术能力：高带宽、低时延的智算互联架构；多层级互联兼容性与可扩展性；弹性资源调度与成本控制^[17]。

1.3 边缘AI场景

边缘AI^[18]定义为AI计算在网络边缘、靠近数据的位置，而非在云端完成，优势在于降低将数据发送到云端的成本、保护敏感数据、实时处理数据以及减少对网络的依赖等^[19]。边缘AI应用要求智算互联具备：低至毫秒级的时延；高带

宽，满足视频类应用需求；高可靠，支持工业类边缘AI应用；支持移动性和异构网络。典型应用如基于5G和云端推理的智慧电厂安全检测与处理。

2 主流厂商解决方案

在智算互联技术快速演进的背景下，主流互联网厂商通过方案创新不断突破传统网络架构的瓶颈，为人工智能、云计算和大规模分布式计算提供核心支撑。字节跳动、阿里巴巴、腾讯和Meta等企业基于自身业务需求推出的高性能网络架构，呈现出异构计算驱动网络重构、软硬协同优化、开放生态融合的技术趋势，为行业提供了具有标杆意义的实践路径。

字节跳动的Megascle集群^[20]通过3层CLOS架构集成超万张GPU。它基于Broadcom Tomahawk5芯片自研网络交换机，Tomahawk5单芯片可提供51.2 Tbit/s带宽和64个800 Gbit/s端口。集群网络设计维持1:1收敛比，服务器配置8张400 Gbit/s网卡直连8个ToR（top of rack）交换机，通过Spine-block和Core-pod分层扩展，结合Swift^[21]+DCQCN^[22]混合拥塞控制算法，显著降低ECMP哈希冲突和网络时延。

阿里巴巴的HPN7网络^[23]采用创新的双上联+多轨+双平面架构，确保网络在超高负载下依然保持高效、稳定的运行，从而满足AI大模型对计算资源的高需求。单Pod支持15 000张GPU。Pod内细分为15个segment，每个segment整合了1 024张主GPU与64张备用GPU，确保了高可用性和卓越的扩展性。自研Solar-RDMA协议栈^[24]以及ACCL通信库^[25]进一步提升网络性能，实现网络的高性能和高稳定互联。

腾讯星脉2.0^[26]基于3级Fat-Tree架构，单集群支持超100 000张GPU。通过整体结构分为Block-Pod-Cluster 3层，其中Block为基础单元，包含256个GPU；Pod以Block为单位进行扩展，

每 Pod 支持 15~64 个 Block, 灵活适应不同的计算需求; Cluster 是最高层级, 最大可支持 16 个 Pod, 以提供强大的算力和扩展能力。网络协议层面自研 TiTa 2.0 协议, 将拥塞控制迁移至端侧网卡, 结合主动拥塞算法以及拥堵智能调度机制, 使通信效率较传统有明显提升。

Meta 的集群组网方案^[8,27]基于 OCP 开放生态构建 Rack-Cluster-Aggregation 三级网络, 单集群可扩展至数万个计算节点。针对 LLM 分布式训练场景, Meta 提出增强型 ECMP (E-ECMP) 方案, 结合 NCCL 库的队列对扩展 (QP scaling) 技术, 对 RoCE 数据包目标 QP 字段进行二次哈希处理, 将 AllReduce 性能提升 40%。在拥塞控制方面, 400 Gbit/s 网络部署中采用接收方驱动流量准入机制, 通过 CTS 数据包同步发送方请求, 动态限制 in-flight 流量规模, 避免传统 DCQCN 算法的性能衰减。同时, 调度系统引入“最小切割”优化策略, 通过感知 GPU 拓扑位置减少跨 Zone 流量。该方案支持 Meta 完成 Llama 405 B 模型的训练。

这些组网方案虽技术路径各异, 但均体现出智算时代网络架构的演进方向: 通过协议栈重构突破传统方案固有局限, 借助可编程数据平面实现业务感知的网络服务, 以及构建软硬协同的端到端优化体系。随着智算集群规模持续扩大, 网络架构的创新将成为突破算力效能天花板的关键变量。

3 智算互联的关键技术

智算互联的实现依赖于从上层协议到底层物理链路的高效协同, 其核心目标在于实现大规模异构算力节点间通信的超低时延、高吞吐量及弹性扩展。

3.1 互联协议

互联协议通过定义标准化的数据交互规则和格式, 确保异构算力节点 (如中央处理器 (cen-

tral processing unit, CPU)、GPU、神经网络处理器 (neural processing unit, NPU)) 间的高效互操作与数据一致性传输; 协议需要在提升算力集群扩展性的同时, 保障复杂计算负载下的通信效率与系统灵活性。

智算集群的不同场景对于互联提出了不同的要求, 在设计中也需要使用不同的应对方式。在智算中心中, 通常使用总线协议进行 Scale Up 扩展支持 TP^[28]以及 1 K 以内的节点扩展, 而使用网络协议进行 Scale Out 扩展支持 DP^[29]、PP^[30]以及 10 K 级别的节点扩展。

总线协议基于内存语义, 通过硬件原生支持的地址空间共享与访存操作, 实现异构计算单元间的无缝协同, 满足实时计算场景对时延与带宽的严苛需求。主流的总线协议包括 PCIe^[31]、CXL^[32]、NVLink^[33]以及最近即将发布的 UALink^[34]。总线协议具有低时延、高效率和大带宽的特点。

除了总线协议以外, 基于消息语义的网络协议也往往用于智算互联。目前主流的网络协议主要有 IB (InfiniBand)^[35]与 RoCE^[36]两大类。IB 网络最初为高性能计算 (HPC) 设计, 具备低时延、高带宽、标准化拓扑管理、拓扑组网方式丰富以及转发效率高的优势, 但存在产业链封闭的问题。RoCE 为数据中心网络设计, 具备高带宽/高弹性组网, 对云化服务支持较好以及生态开放的优点。但是在网络性能方面不如 IB。最近新出现的 UEC^[37]在原有的以太网基础上吸收了 IB 以及总线的一些优点, 性能上有了一定提升, 但是也带来了兼容性方面的问题。

3.2 网络拓扑

智算互联的网络拓扑是实现大规模智算集群的核心。拓扑决定了集群的规模以及带宽、时延等性能指标。Fat-Tree 拓扑^[38]凭借其对称多级交换结构, 在扩展性、时延与带宽间取得平衡。它的无阻塞转发确保任意节点对间均可实现全带宽



通信。此外在数据中心的长期部署使得它是一种非常成熟的大规模集群拓扑，不少智算中心都采用了 Fat-Tree 拓扑，如阿里巴巴的 HPN^[23]、字节跳动的 MegaScale^[20]等。Fat-Tree 可以通过多路径负载均衡技术（如 ECMP）将参数同步流量均匀分散至各层级链路，避免 AllReduce 操作引发网络热点；但其硬件成本与功耗随层级深度呈指数级增长，制约了超大规模集群的部署经济性。

谷歌 TPU^[5-7]和特斯拉 Dojo^[39]使用的是基于两维或多维网格互联结构（如 2D/3D 环形或网格）：mesh^[40]或 torus^[41]拓扑。节点仅与相邻邻居直连，显著降低布线复杂度。此外它还避免了使用商用交换机数量，大大降低了成本。高度定制化特性适配智算互联的专门场景。但是在此类拓扑中全局通信（如广播、归约）因多跳转发所产生的较大累积时延，需借助底层硬件或调度算法解决。

近些年学术界也提出了类似 Hamming Mesh^[42]、TopoOpt^[43]等新颖的拓扑形式，但是尚未有规模的实际商业部署。

3.3 拥塞控制

拥塞控制算法通过动态调节流量注入速率来缓解网络热点对通信时延与吞吐量的影响，它对于提升智算互联的时延和吞吐量指标至关重要。其核心机制在于持续监测网络状态指标（如时延、丢包率、ECN 标记等），并以此为依据调整拥塞窗口或数据包发送间隔。根据感知指标的不同，主流算法可分为以下 4 类。

（1）基于丢包型

以 Reno^[44]、BIC^[45]、Cubic^[46]为代表的算法通过检测数据包丢失触发拥塞动作。这类方案实现简单且部署广泛，但在高速网络场景下存在显著缺陷：丢包事件往往滞后于实际拥塞发生，导致响应速度不足；丢包导致的重传影响网络吞吐量造成带宽利用率剧烈波动。

（2）基于时延型

Swift^[21]、Timely^[47]等算法基于往返时延

（RTT）计算时延梯度（如 RTT 变化率），实现对网络拥塞的早期预测。得益于时延信息的多比特特性，该类算法优势在于收敛稳定且公平性优异，但时延测量易受链路噪声干扰且获取速度较慢，导致算法的反应速度不够理想。

（3）基于 ECN 型

DCTCP^[48]、HPCC^[49]、DCQCN^[22]等算法通过解析交换机标记的显式拥塞通知（explicit congestion notification, ECN）信号进行速率控制。ECN 机制可提供亚 RTT 级别的拥塞感知能力，但单比特标记无法传递拥塞强度信息，导致算法难以精准量化拥塞程度，在异构流量共存场景下易出现振荡现象。此外此类算法的配置数据难以调节也是一个较大的问题。

（4）混合感知型

为突破单一指标的局限性，FASTFLOW^[10]、STrack^[50]等新型算法提出了融合时延梯度与 ECN 标记的多维度反馈机制。它们通过综合使用时延和 ECN 信号，同时提升算法的反应速度和控制精度，取得了良好的效果。

3.4 负载均衡

实际组网中相对的源目节点对之间往往存在多条路径，通过负载均衡算法合理、充分地利用多条潜在路径是解决智算互联中大流量、高突发问题的重要手段，也是提升整体吞吐量的必要选择。

ECMP（equal cost multi-path）^[51]是一种广泛使用的负载均衡算法，它的基本原理是使用哈希函数随机选择一条可用路径来转发特定数据包，该哈希函数通常以数据包头部的特定字段为输入，如由协议号、源地址、目的地址、源端口和目的端口组成的五元组。近年来，又有研究提出在哈希计算中引入额外字段，如生存时间（TTL）或流标签（flow label），以进一步优化路径选择^[52-53]。由于 Hash 函数的静态性，当网络存在拥塞或故障时，多个流依然可能被分配到存在问题的同一条

路径,从而导致更加严重的网络拥塞,甚至丢包。这种问题被称为哈希冲突,它是标准 ECMP 机制的一个较大缺陷。

针对哈希冲突的问题, MPTCP、FlowBender、Flowlet Switching 和 Flowcell 等通过将流拆分为子流(subflow)或流片(flowlet),并分别进行独立路由^[45,54-56]。然而,这些方案仍然无法有效感知网络故障,容易受哈希冲突的影响。无感数据包喷洒(oblivious packet spraying, OPS)^[57]是一种通过在发送端为每个数据包随机选择熵值,或在交换机中随机选择输出端口的技术,将数据包均匀分布到所有可用路径上的技术。OPS 的随机性能够更加均衡分布网络流量,从而缓解哈希冲突。然而,该方法依然无法感知网络拓扑的不对称性或实际流量情况,仍可能无法达到最优负载均衡效果。最新被提出并被 UEC 采用的 Recycled REPS(recycled entropy packet spraying)^[11]是基于熵值实时缓存并筛选优质路径为负载均衡算法,它为负载均衡算法的发展提供了一条新思路。

3.5 通信库

在智算互联生态中,通信库作为分布式算力集群的核心支撑组件,通过优化数据传输协议与硬件资源协同能力,显著提升了大规模模型训练、边缘协同推理等高通信强度场景的效率。当前主流通信库可根据功能以及架构分为 3 类:集合通信库、分区全局地址空间(PGAS)库以及异构通信框架。

集合通信库以多节点集合操作为优化核心,典型代表包括 NVIDIA 开发的 NCCL(NVIDIA collective communications library)^[58]和 AMD 推出的 RCCL(ROCm communication collectives library)^[59]。NCCL 基于拓扑感知算法(如环形/树形通信路径优化)以及底层通信原语的优化,在 NVIDIA GPU 集群可实现高效的 AllReduce、Broadcast 等操作。RCCL 则针对 AMD GPU 架构

优化,兼容 NCCL 接口以降低生态迁移成本。

PGAS 库基于对称性访存模型和单边操作提供分布式系统的访存通信支持,典型实现包括 NVIDIA 提供的 NVSHMEM^[60]及 AMD 提供的 ROC_SHMEM^[61]。

在异构通信框架领域,UCX(unified communication X)^[62-63]对多种网络 API、编程模型、协议和实现进行了抽象封装。其核心思想是提供一组高层通信原语,同时隐藏底层实现的复杂性,由 UCX 运行时负责处理底层通信细节。UCX API 使开发者能够避免依赖于特定厂商的底层驱动,如 cuda 或 ROCm,从而更便捷地提供具有 GPU 感知能力的更加通用的通信功能,但是研究表明这可能以牺牲性能为代价。

尽管通信库技术持续演进,其仍面临多重挑战。首先,硬件生态碎片化导致跨厂商设备(如 NVIDIA 与 AMD GPU)的通信库兼容性问题,增加了系统集成复杂度。其次,超大规模集群(如万卡级训练任务)对通信库的拓扑感知能力提出更高要求,现有算法在动态网络环境下易出现扩展性瓶颈。此外,多租户场景下的安全隔离机制尚不完善,难以保障关键任务的通信服务质量。未来通信库的发展需要在异构兼容性、拓扑感知能力、扩展性和安全性等方面实现突破。

3.6 物理层技术

物理层技术作为算力网络的基础承载,直接决定了数据传输带宽、能效与规模化扩展能力。近年来,电互连、光互连及硅光集成技术的协同突破,持续推动物理层向高密度、低功耗与智能化方向演进,为超大规模智算集群与异构算力互联提供了关键支撑。

当前电互连技术仍是短距离高速传输的主流方案,其核心挑战在于突破信道衰减与串扰对速率提升的限制。基于先进封装(如 2.5D/3D 集成的 CoWoS^[64]、EMIB^[65]、Foveros^[66]等)的互连



架构通过缩短走线长度与阻抗优化,显著提升了电气性能。串行器/解串器(SerDes)作为电互连的核心模块,已向224 Gbit/s^[67]速率迈进。

光互连技术凭借高带宽与低损耗特性,成为长距离及数据中心内部骨干网络的首选。硅光集成通过CMOS兼容工艺将激光器、调制器与探测器集成于单一芯片,大幅降低成本与功耗^[68]。近期,Intel推出的新一代基于硅光的OCI芯粒可在最长为100 m的光纤上单向支持64个32 Gbit/s通道^[69]。共封装光学(co-packaged optics, CPO)技术进一步缩短光引擎与专用集成电路(application specific integrated circuit, ASIC)的距离^[70],Broadcom基于CPO技术推出的Tomahawk 5光互连方案使交换机功耗降低30%,每比特能量效率低于1 pJ^[71]。

光电混合互连架构正成为技术融合的新趋势。针对智算场景的异构流量特征,可重构光电交换技术^[72]通过动态分配电层与光层资源,实现计算-存储-网络资源的灵活组合。此外,量子点激光器等新器件的引入,为未来光学单通道速率提升提供了物理层突破路径。

尽管物理层技术持续革新,其仍面临功耗墙挑战(光模块功耗占比超30%)、多物理场耦合复杂性(如热-电-光协同设计)及规模化制造良率低等瓶颈。未来研究方向将聚焦于光电异构集成、光路可靠性提升、工艺良率提升等方面,以支撑智算互联向ZB级带宽与E级算力的持续演进。

4 发展趋势

智算互联的核心诉求为高带宽、低时延、无阻塞、弹性可扩展以及绿色节能。围绕核心诉求,未来智算互联将在协议与生态、光电混合及全光组网、高效制冷与散热等技术方向持续演进。

4.1 协议与生态

过去,主流智算方案多为封闭的“全家桶”式,各厂商标准不统一,给多元异构算力的互联互通带来极大挑战。未来,开放标准的智算互联将成为主流,以促进整个工业的持续深入发展。以UEC^[37]、UALINK^[34]以及国内的OISA^[73]、ETH+^[74]为代表的智算互联协议开放标准不断涌现,工业界和学术界联手努力促成统一标准来解决算力和互联的效率问题。

4.2 光电混合组网与全光组网

企业将构建10 K甚至更多算力芯片的计算集群以支持大语言模型的技术探索与商用。智算互联的带宽也将快速增长。但电交换芯片的容量提升强依赖半导体工艺提升而进展缓慢,会逐渐成为限制智算互联规模增长的瓶颈。智算互联将基于共封装光学和光学交叉技术的光电混合互联^[75]和全光互连技术^[76]进行组网,以满足未来智算集群规模、带宽、时延的需求。

4.3 高效制冷与散热

当前智算中心的算力需求呈爆炸式增长,与之匹配的高端互联交换芯片功耗已经突破千瓦,这不仅对智算设备和数据中心提出了更高的性能要求,同时也对散热技术提出了严峻挑战。风冷氟泵空调相较于传统的风冷系统具有高效节能、安全可靠、部署灵活等优势,适用于中小型智算中心等算力密度要求不高的应用场景。随着智算中心算力密度不断提升,风冷散热系统存在很大的局限性。液冷散热技术(冷板式、浸没式)相较于风冷具有换热系数高、精准散热、节能等优点,全液冷散热逐步成为高密度智算计算和互联设备的首选散热方式^[77-78];但是液冷也面临材料兼容性、流控稳定性、设计可靠性等问题,需要未来着力解决。

5 结束语

智算互联产业在工业竞争和国家政策两条主

线牵引下,正处于快速发展阶段。智算互联作为支撑人工智能发展的关键基础设施,其技术体系正不断完善,从集群内网络的新型协议与架构创新、生态统一与优化、光电融合创新,再到广域联合训练数据传输技术升级,都在朝着满足日益严苛的智算需求方向演进。通过不断创新与优化,智算互联将逐步形成更加丰富、完善的智算服务生态系统。这不仅将为中国数字经济发展注入源源不断的强大动力,推动各行业的数字化转型与创新发展,还将助力中国在全球科技竞争中占据更为重要的地位,引领世界智算互联技术发展的潮流,为构建智能化社会奠定坚实基础。

参考文献:

- [1] WASWANI A, SHAZEER N, PARMAR N, et al. Attention is all you need[EB]. 2017.
- [2] ACHIAM J, ADLER S, AGARWAL S, et al. GPT-4 technical report[EB]. 2023.
- [3] BI X, CHEN D, CHEN G, et al. DeepSeek LLM: scaling open-source language models with longtermism[EB]. 2024.
- [4] VAVRE A, HE E, LIU D, et al. Llama 3 meets MoE: efficient upcycling[EB]. 2024.
- [5] JOUPPI N P, YOON D H, KURIAN G, et al. A domain-specific supercomputer for training deep neural networks[J]. Communications of the ACM, 2020, 63(7): 67-78.
- [6] NORRIE T, PATIL N, YOON D H, et al. Google's training chips revealed: TPUv2 and TPUv3[C]//Proceedings of the 2020 IEEE Hot Chips 32 Symposium (HCS). Piscataway: IEEE Press, 2020: 1-70.
- [7] JOUPPI N, KURIAN G, LI S, et al. TPU v4: an optically reconfigurable supercomputer for machine learning with hardware support for embeddings[C]//Proceedings of the 50th Annual International Symposium on Computer Architecture. New York: ACM Press, 2023: 1-14.
- [8] AARON G, ABHIMANYU D, ABHINAV J, et al. The Llama 3 herd of models[EB]. 2024.
- [9] WANG W, GHOBADI M, SHAKERI K, et al. How to build low-cost networks for large language models (without sacrificing performance)[EB]. 2023.
- [10] BONATO T, KABBANI A, DE SENSI D, et al. FASTFLOW: flexible adaptive congestion control for high-performance datacenters[EB]. 2024.
- [11] TOMMASO B, ABDUL K, AHMAD G, et al. REPS: recycled entropy packet spraying for adaptive load balancing and failure mitigation[EB]. 2024.
- [12] KAPLAN J, MCCANDLISH S, HENIGHAN T, et al. Scaling laws for neural language models[EB]. 2020.
- [13] WEI J, TAY Y, BOMMASANI R, et al. Emergent abilities of large language models[EB]. 2022.
- [14] GUO D, YANG D, ZHANG H, et al. DeepSeek-R1: incentivizing reasoning capability in LLMs via reinforcement learning[EB]. 2025.
- [15] OpenAI. Learning to reason with LLMs[EB]. 2024.
- [16] BAI J Z, BAI S, CHU Y F, et al. Qwen technical report[EB]. 2023.
- [17] 郭亮, 王少鹏, 权伟, 等. 面向大模型的智算网络发展研究[J]. 电信科学, 2024, 40(6): 137-145.
- [18] GUO L, WANG S P, QUAN W, et al. Research on the development of intelligent computing network for large models[J]. Telecommunications Science, 2024, 40(6): 137-145.
- [19] SHI Y M, YANG K, JIANG T, et al. Communication-efficient edge AI: algorithms and systems[J]. IEEE Communications Surveys & Tutorials, 2020, 22(4): 2167-2191.
- [20] 刘霞. 边缘AI新纪元正在到来[J]. 科技日报, 2024(4).
- [21] LIU X. The new era of edge AI is coming[J]. Science and Technology Daily, 2024(4).
- [22] JIANG Z, LIN H, ZHONG Y, et al. MegaScale: scaling large language model training to more than 10 000 GPUs[EB]. 2024.
- [23] KUMAR G, DUKKIPATI N, JANG K, et al. Swift: delay is simple and effective for congestion control in the datacenter[C]//Proceedings of the Annual Conference of the ACM Special Interest Group on Data Communication on the Applications, Technologies, Architectures, and Protocols for Computer Communication. New York: ACM Press, 2020: 514-528.
- [24] ZHU Y B, ERAN H, FIRESTONE D, et al. Congestion control for large-scale RDMA deployments[C]//Proceedings of the 2015 ACM Conference on Special Interest Group on Data Communication. New York: ACM Press, 2015: 523-536.
- [25] QIAN K, XI Y Q, CAO J M, et al. Alibaba HPN: a data center network for large language model training[C]//Proceedings of the ACM SIGCOMM 2024 Conference. New York: ACM Press, 2024: 691-706.
- [26] MIAO R, ZHU L, MA S, et al. From luna to solar: the evolutions of the compute-to-storage networks in Alibaba Cloud[C]//Proceedings of the ACM SIGCOMM 2022 Conference. New York: ACM Press, 2022.
- [27] DONG J B, WANG S C, FENG F, et al. ACCL: architecting highly scalable distributed training systems with highly effi-



- cient collective communication library[J]. IEEE Micro, 2021, 41(5): 85-92.
- [26] SI J. Development trends of large models and Tencent's independent innovation practice[J]. Bulletin of Chinese Academy of Sciences, 2024, 39(9): 1631-1638.
- [27] GANGIDI A, MIAO R, ZHENG S B, et al. RDMA over Ethernet for distributed training at meta scale[C]//Proceedings of the ACM SIGCOMM 2024 Conference. New York: ACM Press, 2024: 57-70.
- [28] SHOEYBI M, PATWARY M, PURI R, et al. Megatron-LM: training multibillion parameter language models using model parallelism[EB]. 2020.
- [29] VALIANT L G. A bridging model for parallel computation[J]. Communications of the ACM, 1990, 33(8): 103-111.
- [30] HUANG Y P, CHENG Y L, BAPNA A, et al. GPipe: efficient training of giant neural networks using pipeline parallelism[J]. ArXiv e-Prints, 2018, arXiv: 1811.06965.
- [31] LAWLEY J. Understanding performance of PCI express systems[R]. 2014.
- [32] DAS SHARMA D, BLANKENSHIP R, BERGER D. An introduction to the compute express link (CXL) interconnect[J]. ACM Computing Surveys, 2024, 56(11): 1-37.
- [33] FOLEY D, DANSKIN J. Ultra-performance pascal GPU and NVLink interconnect[J]. IEEE Micro, 2017, 37(2): 7-17.
- [34] ARSID R. Ultra Ethernet and UALink: next-generation interconnects for AI infrastructure[J]. IJSAT-International Journal on Science and Technology, 2025, 16(2): IJSAT25023103.
- [35] HAI J, RAJKUMAR B TONI C. An introduction to the InfiniBand architecture[M]//High Performance Mass Storage and Parallel I/O. Piscataway: IEEE, 2009: 617-632.
- [36] ZHU Y, ERAN H, FIRESTONE D, et al. Congestion control for large-scale RDMA deployments[J]. ACM SIGCOMM Computer Communication Review, 2015, 45(4): 523-536.
- [37] METZ J. Empowering AI workloads in ultra Ethernet consortium[C]//Proceedings of the 2024 IEEE Photonics Society Summer Topicals Meeting Series (SUM). Piscataway: IEEE Press, 2024: 1-2.
- [38] LEISERSON C E. Fat-Trees: universal networks for hardware-efficient supercomputing[J]. IEEE Transactions on Computers, 1985, C-34(10): 892-901.
- [39] TALPES E, WILLIAMS D, DAS SARMA D. Dojo: the micro-architecture of tesla's exa-scale computer[C]//Proceedings of the 2022 IEEE Hot Chips 34 Symposium (HCS). Piscataway: IEEE Press, 2022: 1-28.
- [40] DUATO J, YALAMANCHILI S, NI L. Interconnection networks—an engineering approach[M]. Burlington: Morgan Kaufmann, 2002.
- [41] DALLY W J, TOWLES B. Route packets, not wires: on-chip interconnection networks[C]//Proceedings of the 38th Design Automation Conference. Piscataway: IEEE Press, 2005: 684-689.
- [42] HOEFLER T, BONATO T, DE SENSI D, et al. Hamming-Mesh: a network topology for large-scale deep learning[C]//Proceedings of the SC22: International Conference for High Performance Computing, Networking, Storage and Analysis. Piscataway: IEEE Press, 2022: 1-18.
- [43] WANG W Y, KHAZRAEE M, ZHONG Z Z, et al. TopoOpt: co-optimizing network topology and parallelization strategy for distributed training jobs[EB]. 2022.
- [44] PADHYE J, FIROIU V, TOWSLEY D F, et al. Modeling TCP Reno performance: a simple model and its empirical validation[J]. IEEE/ACM Transactions on Networking, 2000, 8(2): 133-145.
- [45] XU L S, HARFOUSH K, RHEE I. Binary increase congestion control (BIC) for fast long-distance networks[C]//Proceedings of the IEEE INFOCOM 2004. Piscataway: IEEE Press, 2004: 2514-2524.
- [46] HA S, RHEE I, XU L S. Cubic: a new TCP-friendly high-speed TCP variant[J]. ACM SIGOPS Oper Syst Rev, 2005, 42: 64-74.
- [47] MITTAL R, LAM V T, DUKKIPATI N, et al. Timely: RTT-based congestion control for the datacenter[C]//Proceedings of the 2015 ACM Conference on Special Interest Group on Data Communication. New York: ACM Press, 2015: 537-550.
- [48] ALIZADEH M, GREENBERG A, MALTZ D, et al. DCTCP: efficient packet transport for the commoditized data center[C]//Proceedings of the ACM SIGCOMM 2010 Conference. New York: ACM Press, 2010.
- [49] LI Y L, MIAO R, LIU H H, et al. HPCC: high precision congestion control[C]//Proceedings of the ACM Special Interest Group on Data Communication. New York: ACM Press, 2019: 44-58.
- [50] LE Y F, PAN R, NEWMAN P, et al. STrack: a reliable multipath transport for AI/ML clusters[EB]. 2024.
- [51] HOPPS C. Analysis of an equal-cost multi-path algorithm[EB]. 2000.
- [52] KABBANI A, VAMANAN B, HASAN J, et al. FlowBender: flow-level adaptive routing for improved latency and throughput in datacenter networks[C]//Proceedings of the 10th ACM International on Conference on Emerging Networking Experiments and Technologies. New York: ACM Press, 2014: 149-160.
- [53] QURESHI M A, CHENG Y, YIN Q W, et al. PLB: congestion signals are simple and effective for network load balancing[C]//

- Proceedings of the ACM SIGCOMM 2022 Conference. New York: ACM Press, 2022: 207-218.
- [54] HE K Q, ROZNER E, AGARWAL K, et al. Presto: edge-based load balancing for fast datacenter networks[C]//Proceedings of the 2015 ACM Conference on Special Interest Group on Data Communication. New York: ACM Press, 2015: 465-478.
- [55] RAICIU C, BARRE S, PLUNTKE C, et al. Improving datacenter performance and robustness with multipath TCP[J]. ACM SIGCOMM Computer Communication Review, 2011, 41(4): 266-277.
- [56] ERICO V, RONG P, MOHAMMAD A, et al. Let it flow: resilient asymmetric load balancing with flowlet switching[C]//Proceedings of the 14th USENIX Symposium on Networked Systems Design and Implementation (NSDI 17). Boston: USENIX Association, 2017: 407 - 420.
- [57] DIXIT A, PRAKASH P, HU Y C, et al. On the impact of packet spraying in data center networks[C]//Proceedings of the 2013 IEEE INFOCOM. Piscataway: IEEE Press, 2013: 2130-2138.
- [58] NAINENI N, JAIN S. State-of-the-Art MPI AllReduce implementations for distributed machine learning: a survey[J]. 2024.
- [59] HE S M, WAN W, LI J H. Network communication optimization of RCCL communication library in Multi-NIC systems[C]//Proceedings of the Third International Conference on Algorithms, Microchips, and Network Applications (AMNA 2024). Bellingham: SPIE, 2024: 28.
- [60] HSU C H, IMAM N, LANGER A, et al. An initial assessment of NVSHMEM for high performance computing[C]//Proceedings of the 2020 IEEE International Parallel and Distributed Processing Symposium Workshops (IPDPSW). Piscataway: IEEE Press, 2020: 1-10.
- [61] PUNNIYAMURTHY K, HAMIDOUCE K, BECKMANN B M. Optimizing distributed ML communication with fused computation-collective operations[C]//Proceedings of the SC24: International Conference for High Performance Computing, Networking, Storage and Analysis. Piscataway: IEEE Press, 2024: 1-17.
- [62] SHAMIS P, VENKATA M G, LOPEZ M G, et al. UCX: an open source framework for HPC network APIs and beyond[C]//Proceedings of the 2015 IEEE 23rd Annual Symposium on High-Performance Interconnects. Piscataway: IEEE Press, 2015: 40-43.
- [63] UNAT D, TURIMBETOV I, ISSA M K T, et al. The landscape of GPU-centric communication[J]. arXiv preprint arXiv: 2409.09874, 2024.
- [64] HOU S Y, CHEN W C, HU C, et al. Wafer-level integration of an advanced logic-memory system through the second-generation CoWoS technology[J]. IEEE Transactions on Electron Devices, 2017, 64(10): 4071-4077.
- [65] CHO J, et al. Electrical characterization of embedded multilayer interconnect bridge (EMIB) and interposer considering system bandwidth and I/O power consumption[C]//Proceedings of the 2017 DesignCon. [S.l.:s.n.], 2017.
- [66] INGERLY D B, AMIN S, ARYASOMAYAJULA L, et al. Foveos: 3D integration and the use of face-to-face chip stacking for logic devices[C]//Proceedings of the 2019 IEEE International Electron Devices Meeting (IEDM). Piscataway: IEEE Press, 2019: 19.6.1-19.6.4.
- [67] KIM J, KUNDU S, BALANKUTTY A, et al. A 224-Gb/s DAC-based PAM-4 quarter-rate transmitter with 8-tap FFE in 10-nm FinFET[J]. IEEE Journal of Solid-State Circuits, 2022, 57(1): 6-20.
- [68] KHANI M, GHOBADI M, ALIZADEH M, et al. SiP-ML: high-bandwidth optical network interconnects for machine learning training[C]//Proceedings of the 2021 ACM SIGCOMM 2021 Conference. New York: ACM Press, 2021: 657-675.
- [69] FATHOLOLOUMI S, MALOUIN C, HUI D, et al. Highly integrated 4 Tbps silicon photonic IC for compute fabric connectivity[C]//Proceedings of the 2022 IEEE Symposium on High-Performance Interconnects (HOTI). Piscataway: IEEE Press, 2022: 1-4.
- [70] MOAZENI S. Next-generation co-packaged optics for future disaggregated AI systems[EB]. 2022.
- [71] Broadcom Ships Tomahawk 5. Industry's highest bandwidth switch chip to accelerate AI/ML workloads[EB]. 2022.
- [72] GUO X T, XUE X W, YAN F L, et al. DACON: a reconfigurable application-centric optical network for disaggregated data center infrastructures[J]. Journal of Optical Communications and Networking, 2022, 14(1): A69.
- [73] 李莹, 王升, 张昊. 打造算网智一体化 AI-Native 算力网络 推动全国一体化算力网纵深发展[J]. 通信世界, 2025(11): 22-23.
- LI Y, WANG S, ZHANG H. Building an AI-native computing network with integrated computing network and intelligence, and promoting the deep development of the national integrated computing network[J]. Communications World, 2025(11): 22-23.
- [74] 厉俊男, 李韬, 杨惠. 超以太网技术的现状与展望[J]. 中兴通讯技术, 2024, 30(6): 48-53.
- LI J N, LI T, YANG H. Status and prospect of ultra-Ethernet technology[J]. ZTE Technology Journal, 2024, 30(6): 48-53.
- [75] HAN X C, ZHAO S Z, LYU Y X, et al. LumosCore: highly scalable LLM clusters with optical interconnect[EB]. 2024.



- [76] WU Z G, DAI L Y, NOVICK A, et al. Peta-scale embedded photonics architecture for distributed deep learning applications[J]. Journal of Lightwave Technology, 2023, 41(12): 3737-3749.
- [77] 李云飞. 关于数据中心液冷技术应用现状和趋势研究[J]. 中国新通信, 2022, 24(12): 72-74.
- LI Y F. Research on application status and trend of liquid cooling technology in data center[J]. China New Telecommunications, 2022, 24(12): 72-74.
- [78] 中国信科. 迈向绿色高效: 数据中心制冷技术的革新之旅[J]. 通信世界, 2024(18): 21-22.
- CICT. Innovative journey towards green and efficient data center refrigeration technology[J]. Communications World, 2024(18): 21-22.

[作者简介]



张云勇 (1976-), 男, 博士, 中国联合网络通信有限公司云南省分公司正高级工程师, 百千万人才工程国家级人选, 国务院政府特殊津贴专家, 云南省“兴滇英才计划”产业创新人才, 主要研究方向为数字经济、人工智能、信息技术、通信技术。



闫硕 (1987-), 男, 博士, 中国联合网络通信有限公司云南省分公司高级工程师, 建设发展部(科技创新部)科技创新管理室主任, 主要研究方向为数理逻辑、形式化方法、人工智能。



陈永铭 (1978-), 男, 现就职于中兴通讯股份有限公司, 主要研究方向为网络架构、片上互联、处理器微架构。



张启明 (1970-), 男, 现就职于中兴通讯股份有限公司, 主要研究方向为智算网络互联架构。