

智算网络专题导读



专题策划人



赵慧玲

工业和信息化部信息通信科技委常委、中国通信学会理事、中国通信标准化协会（CCSA）网络与业务能力技术工作委员会（TC3）主席、中国通信学会标准工作委员会副主任委员、中国电信科技委常委；长期从事电信网络领域的技术与标准工作；曾荣获多项国家及省部级科技进步奖；已发表论文100余篇，出版技术专著12部。

智算网络作为算力高效调度与资源协同联动的核心载体，直接决定了算力资源的利用效率，是支撑AI大模型训练及“东数西算”工程落地的关键基础设施。当前，智算网络面临的核心技术挑战主要集中在4个方面：一是万卡级集群互联扩展性不足，传统网络架构难以支撑高带宽、低时延、无阻塞的大规模组网需求；二是广域无损传输难以保障，远程直接内存访问（RDMA）对丢包极为敏感，在长距场景下，现有拥塞控制与流控机制往往失效；三是大模型流量特征与调度策略不匹配，低熵大象流以及点对多点（P2MP）、全互联（All-to-All）等通信模式易导致负载不均、资源利用率低下；四是存算网边全栈协同缺失，键值缓存（KV-Cache）复用效率偏低，边缘侧资源受限，异构算力与光电混合组网缺乏统一的调度机制。本期专刊特邀互联网和电信运营商、科研院所及高校的专家，聚焦上述核心技术瓶颈，撰写了8篇高质量研究成果。文章既关注技术突破，也注重实践落地，系统梳理了行业最新进展，以期为业界同仁提供切实有益的参考。

针对AI大模型跨园区训练中的长距组网需求，《长距高性能网络关键技术研究》一文聚焦长距传输时延激增、拥塞反馈不及时及端口负载不均等关键问题，提出了一种端到端

联合优化方案。方案采用分级流控技术，以灵活适配数据中心内部流量与跨园区流量的差异化带宽需求；同时设计了基于单跳信令通道的快速拥塞反馈机制，显著降低反馈延时。实验结果表明，该方案将跨园区算力损耗从10%~15%有效控制在5%以内，实测效果显著。

面向广域智算网络中RDMA传输易丢包、拥塞控制滞后等挑战，《广域无损智算网络关键技术研究与实践》提出了一套广域无损智算网络技术方案。方案通过分段反压精准流控与大缓存增强，实现零丢包传输；并引入基于IPv6的段路由（SRv6）切片技术，实现租户级隔离，有效解决优先级流控制（PFC）拥塞扩散问题。文章基于现网实际验证，在云边协同推理和存算分离训练等场景下均能保持高算效，技术路线清晰、工程落地性强，为广域高性能智算组网提供了重要参考。

聚焦于AI大模型产业化落地中的算力互联关键问题，《Scale-Up网络内GPU卡间互联及OISA技术体系》一文系统梳理了Scale-Up卡间互联技术的演进历程与产业趋势，对比分析了私有协议与开放标准的优劣。文章重点介绍了开放互联架构（OISA），从带宽、时延、扩展性及生态兼容性等产业视角，阐述了其在超节点部署、多图形处理器（GPU）协同、大模型训练与推理等场景中的实用价值。该文紧贴产业痛点，兼顾技术落地与生态构建，对智算中心规划、硬件选型与集群建设具有重要指导意义。

针对传统纯电交换网络在带宽、时延、能耗方面存在的

DOI: 10.12142/ZTETJ.202603001
收稿日期: 2026-05-14

瓶颈,《智算中心光电混合组网下的流量差异化调度技术研究》充分利用光电混合组网中“电交换灵活可控、光交换高带宽低时延”的互补优势,构建了“感知—决策—执行—反馈”闭环调度框架。该框架可实现AI训练中数据并行(DP)与流水线并行(PP)流量的精准识别及链路匹配,核心策略是将DP流量优先调度至光链路,而将PP流量动态适配于光或电链路,从而有效提升算力释放效率与链路利用率。

针对边缘大模型推理中显存受限、输入输出(I/O)阻塞、垃圾回收(GC)抖动以及调度粒度过粗等四大瓶颈,《面向边缘场景的高效大模型分层部署机制》提出了Triad-Flow轻量化部署方案。方案采用内存映射(mmap)按需加载技术,破除权重I/O瓶颈;通过静态显存池与指针复用机制,消除运行时抖动;并基于权重矩阵级流水线实现计算与加载的完全重叠。在不增加显存占用的前提下,推理吞吐量提升超过一倍,时延显著降低。

《Token原生的Scale Across网络》一文针对AI时代广域网络互联面临的大规模跨域算力调度挑战,提出以Token为核心调度单元,重构跨域算力互联体系。通过模型与基础设施协同设计、每作业流量工程(Per-job TE)细粒度流量调度、统一IPv6及SRv6寻址,并结合大带宽扁平组网、快速拥塞通知(CNP)、报文裁剪以及链路检测与路由检测(LD/RD)快速倒换等机制,实现了容量扩展能力提升100倍、收敛时间缩短至原来的1/10、成本降低50%的显著效果,为跨域训练推理、存算分离及多可用区(AZ)分布式训练提

供了高性能网络底座。

针对智算网络对高吞吐、低时延及无损传输的严苛需求,《高性能智算网络架构探讨与实践》一文提出了一种端网协同、精准流控、无状态组播与轻量级在网计算一体化的架构方案。具体而言,通过端网速率协商与配额调度,实现主动拥塞规避;基于SRv6开展切片级精准流控;采用比特索引显式复制(BIER)无状态组播技术,优化混合专家模型(MoE)中点对多点(P2MP)通信效率;并利用在网聚合降低多点对一点(MP2P)回传开销。原型验证结果表明,该方案显著消除了调速锯齿现象,提升了传输通量,且兼具标准兼容性与工程落地性。

针对分布式大模型推理中KV Cache复用率低、通信冗余及延迟偏高等瓶颈,《面向分布式推理的大模型键值缓存分发网络》一文提出了键值缓存分发网络(KVCDN),将传统内容分发网络(CDN)升级为张量感知、状态管理的智能缓存架构。方案采用端侧语义规范化、边缘分级缓存与核心全局去重三层设计,通过提示词(Prompt)标准化、语义路由、前缀去重及热点预放置等机制,实现高维张量的高效复用。该方案有效降低了冗余计算与跨网传输开销,为云边协同推理提供了高性能底座。

本次专刊收录的研究成果均立足当前实际技术痛点,为智算网络的技术突破与产业落地提供了重要参考。期望以本专刊为载体,促进业界同仁在智算网络领域深化技术交流与合作,并衷心感谢各位作者的辛勤付出与创新贡献。