



专栏：面向大模型的网络技术

面向大模型的智算网络发展研究

郭亮^{1,2}, 王少鹏¹, 权伟², 李洁¹

(1. 中国信息通信研究院云计算与大数据研究所, 北京 100191;

2. 北京交通大学电子信息工程学院, 北京 100044)

摘要：近年来，全球进入智能计算的蓬勃发展期，作为具有巨量参数和复杂结构的深度学习模型，大模型训练需要在多卡、多服务器间实现训练参数的快速同步，所以对算力中心网络的带宽、时延、可靠性、可扩展性和安全性等提出更高要求。研究了面向大模型训练的智算网络的需求和相关关键技术，对智算网络的研究成果、标准规范和案例实践进行了分析，以期进一步促进智算网络的发展。

关键词：大模型；智算中心；网络技术

中图分类号：TP393

文献标志码：A

doi: 10.11959/j.issn.1000-0801.2024147

Research on the development of intelligent computing network for large models

GUO Liang^{1,2}, WANG Shaopeng¹, QUAN Wei², LI Jie¹

1. Cloud Computing and Big Data Research Institute, China Academy of Information and Communications Technology, Beijing 100191, China

2. School of Electronic and Information Engineering, Beijing Jiaotong University, Beijing 100044, China

Abstract: In recent years, the world has entered a period of vigorous development in intelligent computing. As deep learning models with huge parameters and complex structures, large model training requires fast synchronization of training parameters between multiple cards and servers, which imposes higher requirements on the bandwidth, latency, reliability, scalability and security of datacenter networks. The requirements and related key technologies of intelligent computing networks for large model training were studied, and the standard specifications, academic research, and case practices of intelligent computing networks were analyzed, in order to promote the development of intelligent computing networks.

Key words: large model, intelligent computing center, network technology

收稿日期：2024-04-03；修回日期：2024-05-14

通信作者：王少鹏, wangshaopeng@caict.ac.cn

基金项目：新一代人工智能国家科技重大项目 (No.2021ZD0113003)

Foundation Item: The National Science and Technology Major Project (No.2021ZD0113003)



0 引言

算力中心网络一直是学术界和产业界的热门话题。特别是近年来，随着以 ChatGPT 和 Sora 等为代表的大模型技术和产品不断创新突破，计算产业迈进“智能算力”发展新周期，带动产业格局重构、重塑。大模型的应用和人工智能（artificial intelligence, AI）芯片性能的突破产生的巨大计算和存储需求使算力中心网络面临前所未有的挑战。

2018 年和 2021 年，IEEE 分别发布研究报告《The Lossless Network for Data Centers》^[1]和《Intelligent Lossless Data Center Networks》^[2]，就在线数据、深度学习等需求趋势进行了分析，并对数据中心网络面临的高吞吐、低时延，大规模网络中的拥塞控制技术及算法等问题进行了探讨。随着计算和存储技术的发展，网络成为阻碍数据中心算力发挥的最大瓶颈^[3]。大模型的出现，使得对智算网络的需求变得更加明确。

大模型是指包含超大规模参数的深度学习模型，是人工智能发展取得的重要进展，在图像识别和视频处理等领域具有广泛应用。大模型具备“涌现能力”，即当数据量达到一定程度后，能够从原始训练数据中自动学习并发现新的、更高层次的特征和模式。大模型还可以解决跨领域、跨任务的问题，实现多任务学习、迁移学习和领域自适应等功能，提高模型的泛化性和鲁棒性。因此，发展大模型可以解决许多传统机器学习方法难以处理的问题，推动人工智能技术和应用的进步。大模型需要大算力，大算力需要好网络。字节跳动与北京大学联合发表的论文^[4]显示，该公司已经部署了一个万卡集群，在 42 h

内完成了规模庞大的 GPT-3 模型训练，实现了 55.2% 的算力利用率，给智算网络的创新带来了新的实践。

1 大模型带来的新需求

算力中心包括数据中心、智算中心和超算中心等。算力网络，为连接算力中心的网络，包括用户与算力中心之间的网络、算力中心内部的网络、多个算力中心间的网络，即入算网络、算内网络和算间网络。通过入算网络，企业及个人用户可以连接到算力中心，对接算力应用；通过算内网络，实现算力中心内部存算设备连接，可使计算发挥更高的效用；通过算间网络，基于任务调度等方式实现算力资源共享。智算网络一般属于算内网络的范畴。智算网络作为训练集群架构的基础（见表 1），其重要性不言而喻。

大模型对网络带来的需求有多种。首先是超大模型训练，例如，ChatGPT、Gemini 等万亿参数大模型的训练，需要上万或者数万张图形处理单元（graphics processing unit, GPU）卡，网络规模极大；其次是 Llama 等中小规模的模型训练，采用多任务并行训练的方式，对负载均衡、拥塞控制等方面要求较高；最后是大模型推理，应重点关注网络动态时延。

从应用层看，无论是训练还是推理，大模型处理的数据量都极为庞大，数据传输能力成为训练效率提升的关键，传统的网络架构和互联技术较难满足这种需求。面向大模型的智算网络需要应对分布式存储、并行计算等多节点协同高效通信的场景，具有高带宽、低时延、高可靠和高拓展性等典型特征，须部署更先进的自动化网络解决方案，以支持大规模的数据传输和交互，确保

表 1 训练集群架构

AI 服务	深度学习	框架	计算平台	网络	芯片
图像识别、语音识别、自然语言处理、文生视频…	开发、训练、推理	TensorFlow、Pytorch、MXNet、PaddlePaddle…	CUDA、ROCm、DTK、CANN…	RoCE、InfiniBand	AI 芯片

模型训练和推理的高效运行。

从网络层看, 智算网络中计算与计算、计算与存储之间的流量模型较传统模型差异巨大。传统数据中心网络需要处理大量的南北向流量; 而在由数千个GPU组成的智算集群的基础网络中, 东西向流量更为主要, 且具有大带宽、低时延和全天候的特点。随着微服务体系结构的日益普及, 远程过程调用 (remote procedure calls, RPC) 生成了数据中心中的大部分流量。截至2021年, RPC在谷歌数据中心中产生了95%以上的流量, 其中约75%是进出存储系统的流量^[5]。

2 智算网络关键技术

2.1 网络拓扑架构

组网架构决定了网络的拓扑结构和数据流向, 对网络的稳定性和可扩展性具有重要影响。研究显示, 使用吞吐量而不是平分带宽来评估数据中心网络性能可以改变先前工作中关于数据中心成本、可管理性和可靠性的结论^[6]。目前, 业界主要从以下两个方面进行创新。

(1) 在胖树架构^[7]基础上做大网络容量。胖树架构是对传统树状结构的改进, 从叶子节点到根节点均保持较高带宽。胖树通过增加端口数来扩展网络规模, 不需要改变现有拓扑; 胖树提供与服务器数量成正比的总带宽, 保证了任意两台服务器之间都有足够的通信能力, 这种架构与训练模型相对解耦。

(2) 研究针对特定训练工作负载优化的拓扑, 其中Dragonfly^[8]是比较典型的代表。这种拓扑在为大语言模型训练负载构建专用集群的大企业中较常见。该拓扑中, 每个网络平面拥有独立的硬件资源和软件配置, 负责处理一部分特定的数据流, 它们之间通过特定的连接点进行通信和数据交换。这种架构具有高可靠性和负载均衡等优势; 但每次增加网络容量架构都必须重新调整, 提高了网络的复杂性和管理难度。

传统的三层全连接电交换架构已经不能很好地满足智算中心的需求, 需要一种更高效和易于扩容的电光融合交换或光交换方案来逐步代替传统纯电路分组交换的策略^[9]。文献^[9]认为, 光电融合交换增加了光交换功能, 将波长作为调度单元进行大颗粒交换, 具有可重构、低功耗和低时延等特点, 适用于对性能要求较高的智算中心网络。

2.2 机卡互联技术

大模型训练过程中, 参数通过高速互连网络在不同的服务器间进行同步交互, 传输的数据量较大, 网络易出现负载分担不均、整网吞吐下降的问题, 从而影响大模型训练的性能。除了优化网络拓扑结构, 加强多服务器间多卡互联, 还需要进一步优化数据传输协议, 在物理层和传输控制策略上提供全面支持。大模型训练对算力需求巨大, 多机间的通信是影响分布式训练的一个重要指标。基于典型模型建模发现, 在类似GPT-3的千亿参数模型中, 通信的端到端耗时占比为20%; 针对某个万亿参数模型建模发现, 通信的端到端耗时占比急剧上升到约50%^[10]。在TCP/IP网络中, 数据发送方需要将数据进行多次内存复制, 并经过一系列网络协议数据包处理工作, 其服务器间通信时延为毫秒级别。远程直接数据存取 (remote direct memory access, RDMA) 通过应用程序直接读取或写入远程内存, 避免操作系统、协议栈的介入, 从而实现数据更加直接、简单、高效地传输, 大幅减少数据传输过程中所需的时间^[11]。

无限带宽 (InfiniBand, IB) 技术和基于融合以太网的RDMA (RDMA over converged ethernet, RoCE) 是支持RDMA的重要技术路线。在大模型训练过程中, 业界主要基于这两种技术路线进行改进优化, 打造适应大模型训练的网络技术。英伟达是IB技术的主要推动者, 推出了一系列网络产品和解决方案, 为大模型分布式训练中多机互联提供了重要能力。RoCE技术得到了更

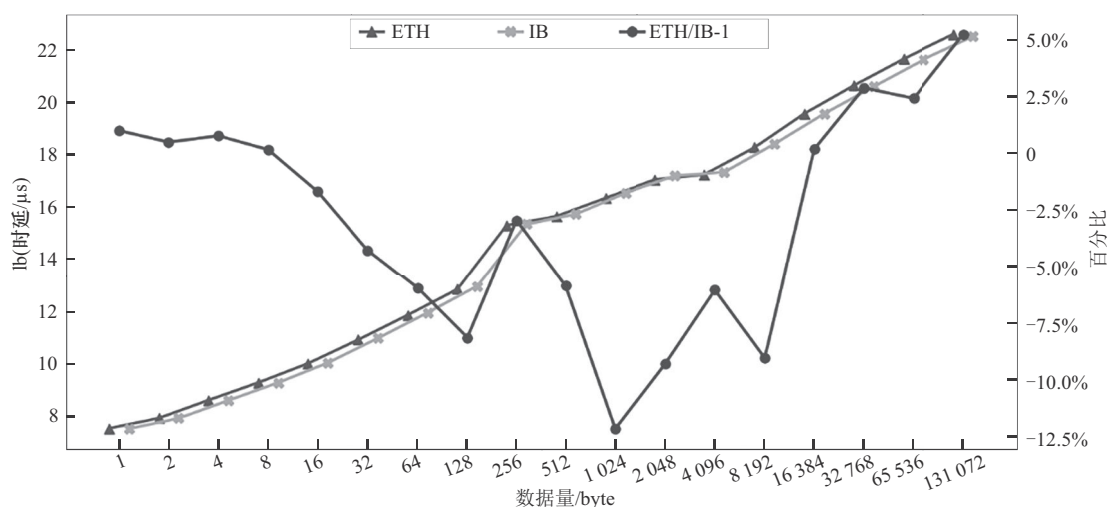


多传统以太厂商的支持。截至2023年6月,在全球超算榜单TOP 100中,IB和RoCE分别占比62%和17%;而在TOP 500中,这一数据变成了40%和45%^[12]。Alltoall测试结果对比和AllReduce测试结果对比分别如图1和图2所示,实测发现,在IMB Benchmark场景下,RoCE的性能不弱于IB,甚至在某些方面还具有一定优势。

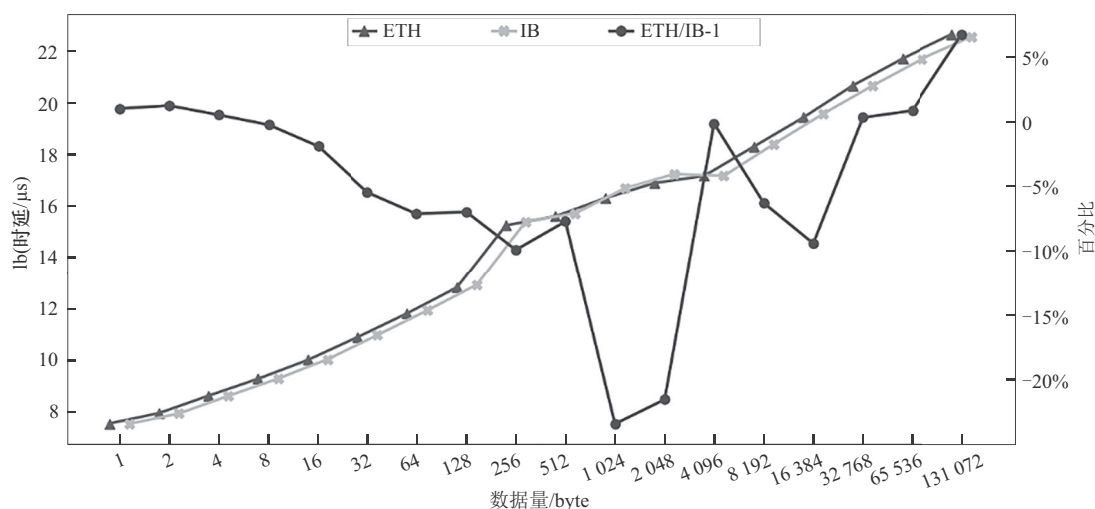
2.3 拥塞控制技术

由于网络流量具有随机性和路径多样性特征,若不采取流量及拥塞控制机制,网络源端数

据发送速率可能超过网络转发传输速率,导致网络拥塞和排队时延增加,影响网络带宽利用率。所以通常采用设置网络拥塞指标或门限,实时感知网络带宽及端口利用率;在出现问题时及时通知源端网卡或设备进行降速,以防止过多数据注入网络。目前,常见的网络拥塞控制技术主要包括4类,分别是基于优先级的流量控制(priority-based flow control, PFC)、显式拥塞通知(explicit congestion notification, ECN)、基于往返时间(round trip time, RTT)的拥塞控制和基于带

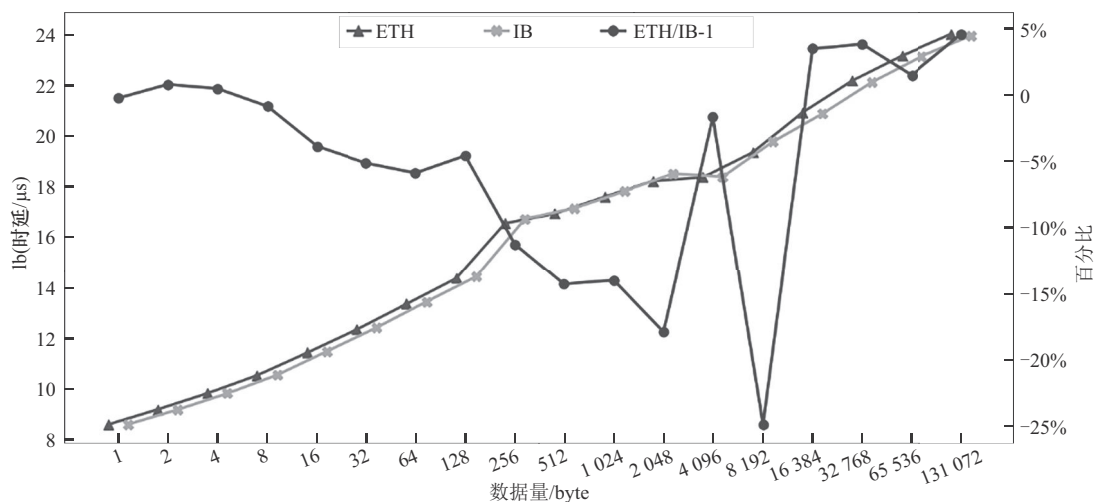


(a) 同 leaf-8 节点

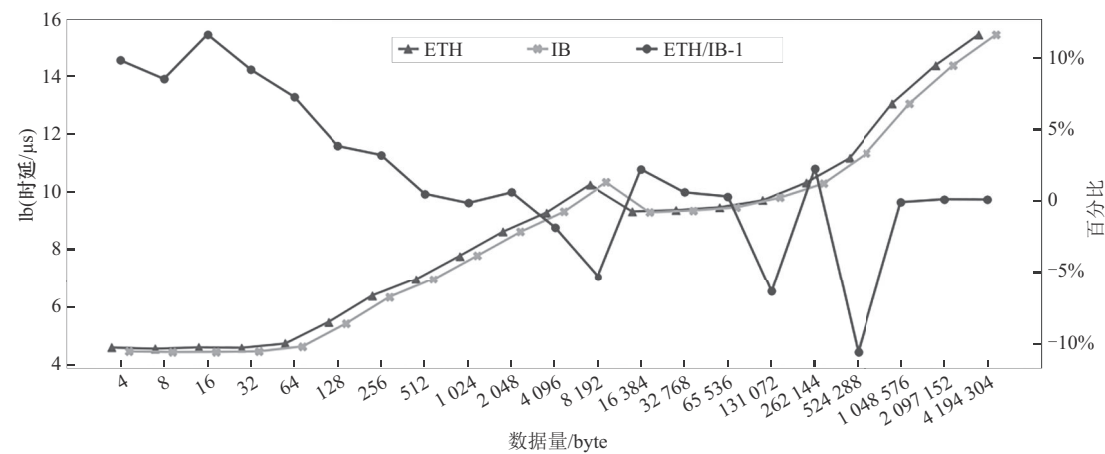
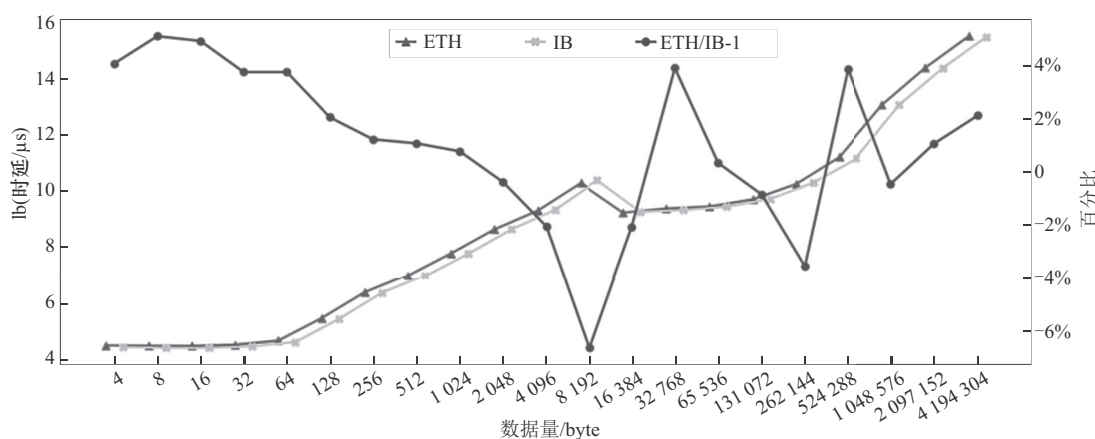


(b) 跨 leaf-8 节点

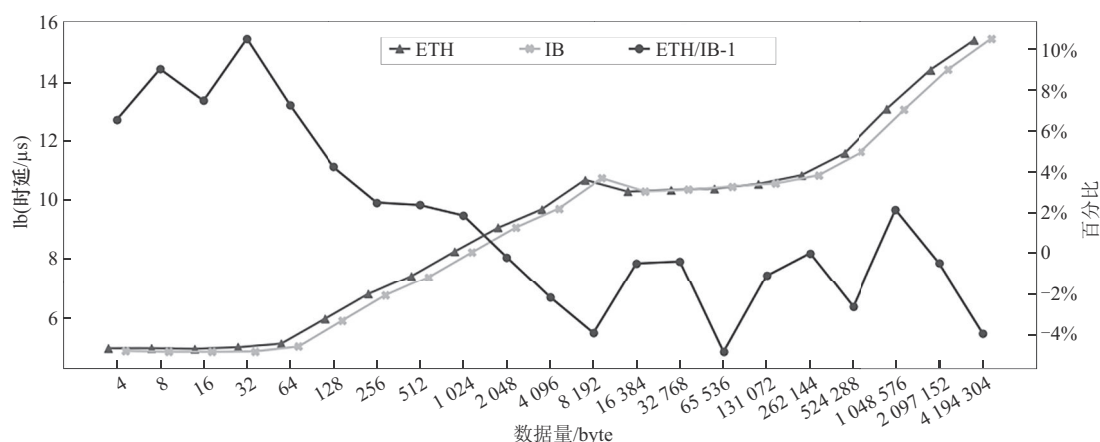
图1 Alltoall测试结果对比



(c) 跨 leaf-16 节点
图1 Alltoall测试结果对比(续)



(b) 跨 leaf-8 节点
图2 AllReduce测试结果对比



(c) 跨leaf-16节点

图2 AllReduce测试结果对比(续)

内网络遥测 (inband network telemetry, INT) 的拥塞控制。

PFC技术主要基于端口队列进行逐级的流量控制,通过在交换机上配置PFC水线,在本地入接口队列发生拥塞或者即将发生拥塞时,向上一级设备发送反压信号,实现对应端口队列停止发送流量,直至解除网络拥塞。

基于ECN的拥塞控制技术是一种以流为粒度的公平性拥塞控制技术。它在设备上配置ECN门限值,当本设备出接口发生拥塞时,将会以一定概率启动ECN标记,接收端将根据ECN标志通知发送端降速。端到端拥塞控制环路失效会导致拥塞扩散,这种非必要拥塞会增加PFC信息和链接暂停时间,也会进一步延迟ECN标记数据包的传输。此时,可以由通知点(notification point, NP)发送拥塞通知包(congestion notification packet, CNP)进行网络智能化补充。只有当出口队列严重拥塞时,才会执行补充CNP操作,因此时延和吞吐量不会受到影响^[13]。

基于RTT的拥塞控制技术主要以主机往返时延衡量拥塞情况,通过将报文在网络中的RTT作为拥塞信号,以RTT梯度调整传输速率,其典型技术方案是谷歌的TIMELY技术。该技术方案对网络设备要求较低,但报文的往返时间易受队列

的影响导致信息滞后而影响网络对拥塞的判断。

阿里云提出了一种基于INT的拥塞控制协议——高精度拥塞控制(high precision congestion control, HPCC)^[14]。相比于数据中心量化拥塞通知(data center quantized congestion notification, DCQCN)算法和Timely算法,HPCC牺牲一定的带宽引入了INT能力,同时也获得了超高精度的拥塞控制性能。HPCC可以实现快速的算法收敛,从而更好地利用闲置带宽,同时保持交换机始终处于近零队列,实现超低的数据传输时延。

2.4 负载均衡技术

负载均衡技术能够有效保障集群内部的高吞吐,尤其是在大模型训练过程中,负载均衡的重要性进一步凸显。随着大模型参数规模的进一步增长,多节点数据同步和传输需求进一步提升,智算中心网络物理带宽逐步从100~200 Gbit/s提升到400~800 Gbit/s,在高带宽基础上实现网络高吞吐,负载均衡将会发挥重要作用。网络均衡的主流技术有3种:逐流等价多路径路由(equal-cost multi-path routing, ECMP)负载均衡、逐包负载均衡和网络级负载均衡(network scale load balance, NSLB)。

智算中心网络通常采用胖树架构,任意出入口端口之间存在多条等价转发路径,逐流ECMP负

载均衡在“流量大、数量少”的场景下易造成哈希极化。无阻塞网络中,链路负载不均衡会导致冲突流,使有效带宽下降以及冲突流序列化时间增加,应用效果不佳。逐包负载均衡机制基于路径的状态信息,针对包进行动态选路,从而达到流量哈希均衡。网络级负载均衡通过绘制全局的流量矩阵计算出最佳的流量图,进行自动导流,实现路径预规划和路由分配,确保数据流量无拥塞,从而达到全网吞吐最优。HCCL环境下开启NSLB前后对比如图3所示,通过实测可见,在包括华为集合通信库(Huawei collective communication library, HCCL)等多种环境下,开启NSLB后,系统整体性能得到明显提升。

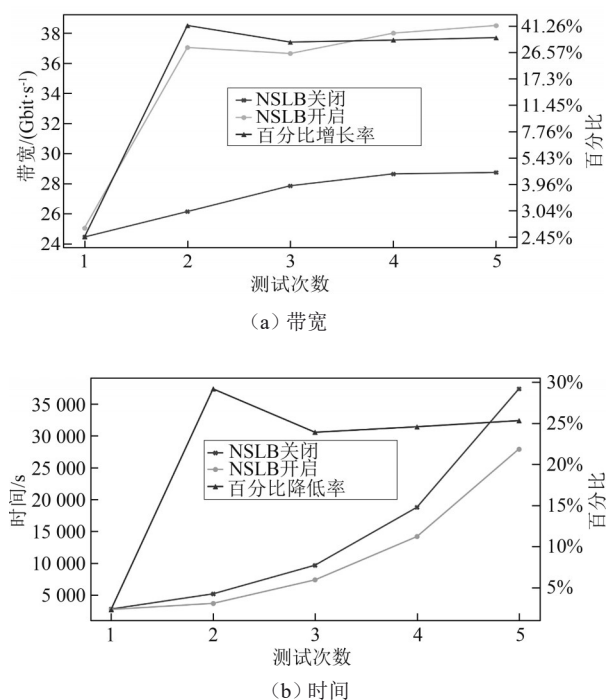


图3 HCCL环境下开启NSLB前后对比

3 智算网络创新实践

大模型的快速兴起加速了高性能网络关键技术的研究和应用的发展,为智算产业发展带来了巨大的发展机遇,云服务提供商、通信设备制造商、电信运营商以及科研机构等多方主体纷纷投

入智算网络技术的研究,并通过跨界融合和深度合作,共同构建起丰富多元、协同创新的产业生态,推动智算网络技术升级。

3.1 成果研究

针对AI训练异构网络问题,亚马逊云科技(Amazon web services, AWS)提出MiCS方案^[15]。MiCS将集群中的GPU划分为多个组,并在每个组内划分模型状态,每个组内的分区模型状态限制在固定数量的GPU内最频繁的通信、参数收集和梯度同步。AWS评估表明,MiCS实现了近线性缩放的效率。微软提出Singularity^[16]以解决AI资源碎片化问题。其核心思路是将一个训练拆分为多个任务,任务之间进行集合通信调度;每个物理GPU被虚拟为一到多个虚拟GPU,每个虚拟GPU承载一个任务;各层调度器在与具体硬件资源解耦的基础上,提供调度功能,实现异构训练卡的池化,提升云上AI资源利用率。

文献[17]提出了融算网络体系,该体系在应用侧,面向各种大规模训练平台、研究数据云、多语义翻译平台等提供差异化的算力支撑与协同调度能力,保障业务的稳定、可靠运行,实现跨节点分布学习,为大模型、大应用等提供超级算力。文献[18]提出了融智算力网络,通过通信网络将可用的计算节点、存储节点连接起来,构成整个体系的硬件底座,然后在其上部署各类智能模型,一方面作为指挥员调度和编排各类资源,为用户提供敏捷精准的按需服务,实现内生智能,另一方面作为智能资源在系统内共享与分发,提升整个系统的业务处理能力。

3.2 标准规范

近几年,中国信息通信研究院联合业界发布了YD/T 3902《数据中心无损网络典型场景技术要求和测试方法》、YD/T 4028《基于RoCE协议的数据中心高速以太无损网络测试方法》、YD/T 4073《基于远程直接内存访问的高速以太存储网络交换设备测试方法》等多个通信行业标



准。开放数据中心委员会的网络工作组和新技术与测试工作组均就智算网络开展了研究，陆续发布了 ODCC-2018-05005《数据中心无损网络总体技术要求》、ODCC-2019-05002《数据中心 AI 及三网合一技术白皮书》、ODCC-2020-05006《开放式拥塞控制 OpenCC 技术白皮书》、ODCC-2021-05009《分布式 AI 训练实践白皮书》、ODCC-2022-05011《下一代以太网技术需求白皮书》、ODCC-2023-0500C《网络级 DMA (NDMA) 技术需求白皮书》等成果。

3.3 案例实践

公开资料显示，华为星河 AI 网络采用全局负载均衡算法、网络智能调优算法和全栈可视运维技术实现算网实时协同调度，将网络有效吞吐从 50% 提升到 98%，大模型训练效率提升 20% 以上；采用冗余链路设计，确保在网络出现故障时能够及时切换，保证大模型训练的连续性。腾讯星脉 AI 网络 2.0 采用自研交换机、端网协同协议 TiTa/CNIC、腾讯集合通信库 (Tencent collective communication library, TCCL) 以及端到端网络运营系统，可提升 40% 的 GPU 利用率，节省 30%~60% 的模型训练成本，为 AI 大模型带来 10 倍通信性能提升^[19]。阿里巴巴的可预期网络，以其接近微秒级的端到端超低时延、千卡规模的每节点 400 Gbit/s 通信带宽，结合阿里巴巴集合通信库 (Alibaba collective communication library, ACCL) 进行优化，可为用户提供更高的通信效率。在百度高性能智算网络方案中，基于导轨优化的大带宽高吞吐的 AI-POOL 智算网络架构、智算网络的多租户方案 AI-VPC、智算网络可视化网管系统 AI-NETOP 都已形成了成熟的产品化能力和解决方案，并落地了商业化的用户案例，可帮助客户高效地运行智算应用^[20]。

4 结束语

除了上述技术以外，自动化部署、网络故障

感知和收敛、数据处理单元 (data processing unit, DPU) 等也是智算网络的重要技术方向。集合上述技术，构建端到端、软硬协同的智算网络体系，并通过协议层的支撑和运维层的管理可为各类 AI 应用提供强大的智算网络能力。

智算在其发展过程中不可避免地会面临一系列挑战。首先是生态尚待完善。中国自主研发的推理芯片技术相对成熟，但基于现场可编程门阵列 (field programmable gate array, FPGA) 和专用集成电路 (application specific integrated circuit, ASIC) 等多种技术路线，面临着对各类开发框架的支持、算法的优化、平台的适配等多个问题，亟须形成协同统一的软硬件生态去推动整个产业的发展。其次是算力合理利用。据不完全统计，截至 2023 年年底，全国与智算中心相关的项目有 120 多个，算力规模从 50 P 到 1 000 P 不等，总量超过 7.7 万 P (P 表示每秒 1 000 万亿次运算)。中小规模算力对大模型训练作用有限，需有强大的“智算大脑”来协同高效利用此类算力。最后是绿色算力发展。通算中心着重关心电源使用效率；到了智算时代，算力碳效等引起了业界更多的关注。绿色算力涉及设施、IT 设备和运营管理等多个层面，网络是其中很重要的一部分。

总体来看，随着大模型应用的逐步深入，全球的智能计算进入快速发展期，由此带来了智能算力的爆发式需求。智能算力的充分释放极大依赖于智算网络的性能、可扩展性、安全性和自动化等多重因素，还需业界进行更多的创新和发展。

参考文献：

- [1] CONGDON P, MARKS R. IEEE 802 Nendica report: the lossless network for data centers[R]. 2018.
- [2] GUO L, CONGDON P. IEEE 802 Nendica report: intelligent lossless data center networks[R]. 2021.
- [3] 孙黎阳, 温小振, 郭亮. 总线级数据中心网络关键技术研究[J]. 中国电信业, 2022(10): 75-77.
SUN L Y, WEN X Z, GUO L. Research on key technologies of bus-level data center network[J]. China Telecommunications Trade, 2022(10): 75-77.

- [4] JIANG Z H, LIN H B, ZHONG Y M, et al. MegaScale: scaling large language model training to more than 10 000 GPUs[J]. arXiv preprint, 2024: arXiv: 2402.15627.
- [5] ZHANG Y W, KUMAR G, DUKKIPATI N, et al. Aequitas: admission control for performance-critical RPCs in datacenters[C]//Proceeding of the 2022 ACM SIGCOMM'22 Conference, New York: ACM, 2022.
- [6] NAMYAR P, SUPITTAYAPORN PONG S, ZHANG M Y, et al. A throughput-centric view of the performance of datacenter topologies[C]//Proceedings of the 2021 ACM SIGCOMM 2021 Conference. New York: ACM, 2021.
- [7] AL-FARES M, LOUKISSAS A, VAHDAT A. A scalable, commodity data center network architecture[C]//Proceedings of the ACM SIGCOMM 2008 Conference on Data Communication. New York: ACM, 2008: 63-74.
- [8] KIM J, DALLY W J, SCOTT S, et al. Technology-driven, highly-scalable dragonfly topology[C]//Proceedings of the 2008 International Symposium on Computer Architecture. Piscataway: IEEE Press, 2008: 77-88.
- [9] 唐雄燕, 魏步征, 沈世奎, 等. 智算数据中心光电交换技术综述[J]. 光通信研究, 2024(4): 1-17.
TANG X Y, WEI B Z, SHEN S K, et al. Overview of optoelectronic switching technology in intelligent computing data centers[J]. Study on Optical Communications, 2024(4): 1-17.
- [10] 何宝宏, 郭亮, 王少鹏. 星河AI网络白皮书[R]. 2023.
HE B H, GUO L, WANG S P. Star river AI network white paper[R]. 2023.
- [11] 王少鹏, 郑常奎, 芦帅, 等. 数据中心无损网络关键技术研究[J]. 信息通信技术与政策, 2021(10): 68-74.
WANG S P, ZHENG C K, LU S, et al. Research on key technologies of data center lossless network[J]. Information and Communications Technology and Policy, 2021(10): 68-74.
- [12] 黄云皓. 大模型时代的“非大模型”机会: 智算中心以太网[R]. 2023.
HUANG Y H. Opportunities for “non large models” in the era of large models: Ethernet in intelligent computing centers[R]. 2023.
- [13] 郭亮, 李洁, 高峰. 数据中心智能无损网络白皮书[R]. 2021.
GUO L, LI J, GAO F. White paper on intelligent lossless networking in data centers[R]. 2021.
- [14] MIAO Y L R, LIU H H. HPCC: high precision congestion control[C]//Proceeding of the 2019 ACM SIGCOMM'19 Conference. New York: ACM, 2019.
- [15] ZHANG Z, ZHENG S, WANG Y D, et al. MiCS: near-linear scaling for training gigantic model on public cloud[J]. ArXiv e-Prints, 2022: arXiv: 2205.00119.
- [16] SHUKLA D, SIVATHANU M, VISWANATHA S, et al. Singularity: planet-scale, preemptive and elastic scheduling of AI workloads[J]. arXiv preprint, 2022, arXiv: 2202.07848.
- [17] 张宏科, 于成晓, 权伟, 等. 融算网络体系基础研究[J]. 电子学报, 2022, 50(12): 2928-2934.
- ZHANG H K, YU C X, QUAN W, et al. Fundamental research on computing integration networking[J]. Acta Electronica Sinica, 2022, 50(12): 2928-2934.
- [18] 胡玉姣, 贾庆民, 孙庆爽, 等. 融智算力网络及其功能架构[J]. 计算机科学, 2022, 49(9): 249-259.
HU Y J, JIA Q M, SUN Q S, et al. Functional architecture to intelligent computing power network[J]. Computer Science, 2022, 49(9): 249-259.
- [19] 腾讯云计算(北京)有限责任公司, 中国信息通信研究院云计算与大数据研究所. 智算赋能算网新应用白皮书[R]. 2023.
Tencent Cloud Computing (Beijing) Co., Ltd., Cloud Computing and Big Data Research Institute of the China Academy of Information and Communication Technology. White paper on new applications of intelligent computing empowerment network[R]. 2023.
- [20] 李兆彤, 史磊, 周磊. 智算中心网络架构白皮书[R]. 2023.
LI Z T, SHI L, ZHOU L. White paper on network architecture of intelligent computing center[R]. 2023.

[作者简介]



郭亮 (1980-), 男, 中国信息通信研究院云计算与大数据研究所总工程师、正高级工程师, 北京交通大学电子信息工程学院博士生, 主要从事与算力相关的政策支撑、技术研究和标准制定工作。



王少鹏 (1992-), 男, 中国信息通信研究院云计算与大数据研究所数据中心部副主任、工程师, 主要从事与算力融合相关的产业咨询、技术研究和标准制定工作。



权伟 (1987-), 男, 博士, 北京交通大学电子信息工程学院教授、博士生导师, 主要从事新型网络体系架构、高可靠网络传输关键技术研究工作。



李洁 (1979-), 女, 博士, 中国信息通信研究院云计算与大数据研究所副所长、正高级工程师, 开放数据中心委员会副主席, 主要从事与算力相关的产业、政策、技术、标准研究工作。