

智算网络发展综述



Development of Intelligent Computing Networks: A Survey

段晓东/DUAN Xiaodong, 程伟强/CHENG Weiqiang,
张昊/ZHANG Hao

(中国移动通信有限公司研究院, 中国 北京 100053)
(The Research Institute of China Mobile, Beijing 100053, China)

DOI: 10.12142/ZTETJ.202502008

网络出版地址: <http://kns.cnki.net/kcms/detail/34.1228.TN.20250430.1509.002.html>

网络出版日期: 2025-04-30

收稿日期: 2025-02-25

摘要: 智算中心网络作为智算中心的连接底座, 需要具备高性能、低时延的通信能力。智算中心网络体系是一个多要素融合的复杂系统, 依赖于智算业务、机内/外交换芯片、网卡、网络设备等上下游产业协同创新。系统剖析了服务器/超节点内图形处理器 (GPU) 互联网络、同园区机间互联网络、跨园区智算中心互联网络三大核心领域, 探讨了智算网络的需求、挑战以及业界发展态势。中国移动创新提出全向智感互联架构 (OISA)、全调度以太网 (GSE) 技术体系、弹性以太网聚合、精细化拥塞控制、物理层安全等多项创新技术, 旨在构建超大规模、超高带宽、超低时延、超高可靠的智算中心网络, 助力人工智能 (AI) 产业发展。

关键词: 智算中心网络; 全向智感互联架构; 全调度以太网; 智算中心互联

Abstract: As the connection base of the intelligent computing center, the network of the intelligent computing center needs to have high-performance and low-latency communication capabilities. The network system of the intelligent computing center is a complex system that integrates multiple elements, relying on collaborative innovation among upstream and downstream industries such as intelligent computing services, forwarding chips, network cards, and network equipment. This paper systematically analyzes three core areas, namely the graphics processing unit (GPU) Internet in servers/super nodes, the Internet between computers in the same park, and the Internet between intelligent computing centers across parks. The requirements, challenges, and development trends in the industry of intelligent computing networks are discussed. China Mobile has innovatively proposed technical systems such as omni-directional intelligent sensing express architecture (OISA) and global scheduling ethernet (GSE), as well as a number of innovative technologies including elastic Ethernet aggregation, refined congestion control, and physical layer security. The goal is to build an intelligent computing center network with super scale, ultra-high bandwidth, ultra-low latency, and ultra-high reliability, so as to boost the development of the artificial intelligence (AI) industry.

Keywords: intelligent computing center network; omni-directional intelligent sensing express architecture; global scheduling ethernet; intelligent computing center interconnection

引用格式: 段晓东, 程伟强, 张昊. 智算网络发展综述 [J]. 中兴通讯技术, 2025, 31(2): 53-62. DOI: 10.12142/ZTETJ.202502008

Citation: DUAN X D, CHENG W Q, ZHANG H. Development of intelligent computing networks: a survey [J]. ZTE technology journal, 2025, 31(2): 53-62. DOI: 10.12142/ZTETJ.202502008

1 AI 业务与智算中心网络的发展

1.1 AI 业务发展趋势

在当今数据驱动的时代, 人工智能 (AI) 大模型正掀起一场前所未有的网络革命, 成为推动智能时代前行的新浪潮。随着 AI 技术的飞速进步, 大语言模型 (LLM) 以其前所未有的能力已迅速成为行业焦点。

大模型遵循三大统计特征: 第一, Scaling Law^[1], 即模型表现依赖于模型规模、计算量和数据量, 这些因素之间呈

现幂律关系; 第二, Chinchilla Scaling Law, 模型大小和数据量要同等比例扩展, 即数据量需达到参数量的 20 倍, 模型训练结果才能达到饱和; 第三, 涌现能力^[2], 产业界发现, 大语言模型相比之前的预训练语言模型 (PLM) 最显著的特征之一是它们的涌现能力。随着模型规模的指数级增长, 模型内部出现了在小规模模型中难以观察到的涌现能力。这种涌现能力使得 AI 执行任务的能力大幅提升, 且该提升与具体模型无关^[3]。这种由量变引起的质变, 不仅体现在处理复杂任务的能力上, 更深刻地改变了我们对智能计算中心 (智算中心) 网络设计和优化的理解。为了不断提高模型的表现能力, AI 大模型发展迅速, 参数量不断增加。以 Chat GPT 为例, 据统计^[4], GPT-1 由 1.17 亿个参数组

基金项目: 国家重点研发计划项目 (2024YFB2906600)

成, GPT-2 中的参数数量增长至 15 亿, GPT-3 中的参数数量为 1 750 亿, 而 GPT-4 则约有 1.8 万亿个参数, 其参数数量约是 GPT-3 的 10 倍, 是 GPT-1 的 15 000 倍。

由于存在算力、内存的限制, 为了提高训练效率和模型泛化能力, 大模型需要通过多张图形处理器 (GPU) 卡进行并行训练。常见的并行训练模式有数据并行 (DP) 和模型并行 (MP), 其中模型并行又分为张量并行 (TP)、序列并行 (SP)、混合专家 (MOE) 并行和流水线并行 (PP)。

1) TP 对模型层内做划分, 将大型矩阵运算 (即张量操作) 分割成更小的部分, 在不同的 GPU 卡上并行执行, 以解决由模型数据导致的内存瓶颈问题。

2) SP 用于解决非模型数据 (如中间特征值) 导致的性能瓶颈问题, 将长序列训练任务分解成多个子序列块并将其保留在不同 GPU 卡中, 使其能够处理更长的输入数据。

3) MOE 并行将大模型拆分成多个小模型 (即专家), 每轮迭代根据输入样本通过门控网络, 分配给一部分专家进行计算, 以提高模型训练效率。

4) PP 将模型划分为多个层, 把不同的层按顺序分配到不同的节点上, 在拆分模型的边界处插入通信步骤, 并在层之间点对点传递中间结果 (包含计算结果的前向传递和梯度的反向传递)。PP 对带宽要求低, 对时延要求不太高。

5) DP 将大规模的数据集划分为多个子集并按批次分配给不同节点。每个节点都独立对各自的子集进行相同的计算

操作, 然后在各个节点之间汇总得到最终结果。GPU 卡间需传输大量梯度数据, 对带宽要求较高。由于通信可以被计算掩盖, 因此 DP 对于时延的要求相对较低。

在实际训练中, 可结合 DP、MP 的优势, 形成混合并行策略。例如, 除了在单机内使用 TP 和 SP 组合的策略外, 还可使用 PP 策略实现跨多台机器计算, 最后通过 DP 增加并发数量。

1.2 智算网络发展趋势

大模型的快速迭代演进和算力资源需求的快速增长, 驱动着智算中心组网规模向万卡级甚至 10 万卡级演进。支撑这类场景的网络通常属于高性能网络范畴, 可同时满足超 10 万节点的组网规模、太比特每秒以上的有效带宽、零丢包、亚微秒级时延和抖动、超高稳定性等极致性能要求。如此大规模的 GPU 互联互通网络共分为 3 个部分: 同服务器/超节点内的 GPU 卡间互联网络、同园区机间互联网络、跨园区智算中心间互联网络, 如图 1 所示。这 3 层网络在延迟、带宽、可靠性方面都存在较大差异。大模型根据这 3 层网络的能力进行模型切分并完成分层交互。每一层网络都有不同的性能要求和挑战, 需要分别进行优化。

1) 同服务器/超节点内的 GPU 卡间互联网络主要解决单服务器内部或超节点内部 GPU 之间互联互通问题, 承载 TP、SP、MOE 并行参数同步等紧密协作任务。TP、SP、MOE 并行时, GPU 卡间因需同步大量参数数据, 且通信和计算不可以掩盖, 对互联网络有纳秒级超低延迟、太比特每秒级超高

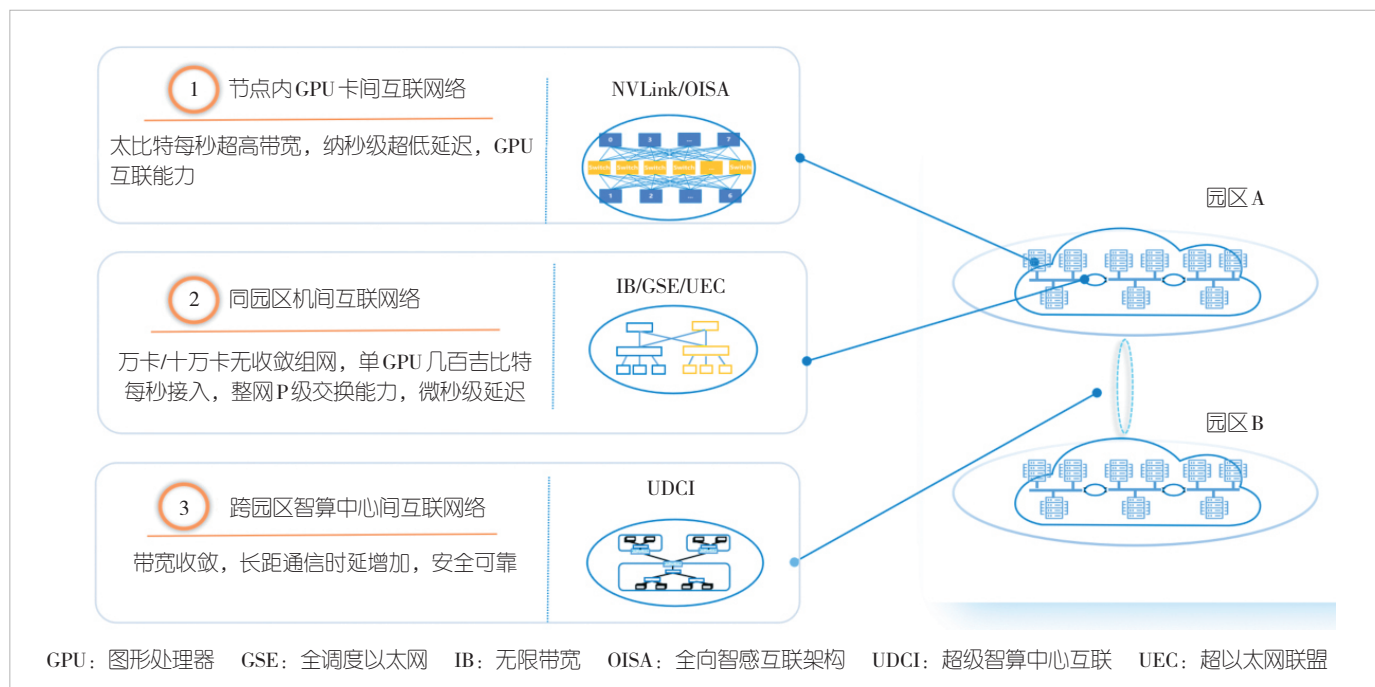


图1 新型智算中心网络

带宽和强一致性需求。机内互联网络的性能直接影响TP、SP、MoE的并行效率，其硬件实现通常依赖专用高速接口，如NVLink或外围组件互连高速总线（PCIe）通道。当前中国智算服务器只支持8卡PCIe直连拓扑、点到点带宽小于64 GB/s，无法满足MoE并行高频率、大带宽、低延时的通信需求。为此，产业界正在积极开展超节点服务器的研究。

2) 同园区机间互联网络主要解决同一智算中心内的跨服务器/机架通信问题，承载PP、DP并行参数或梯度同步任务。一个千亿级大模型训练需要万卡规模的GPU协同、单GPU几百Gbit/s的高带宽接入，并通过远程直接内存访问（RDMA）来减少传输时延。机间互联网络需要具备可规模扩展并支持万卡甚至十万卡无收敛组网。整个智算中心网络支持P级别交换能力和微秒级低延迟，通常使用IB和基于融合以太网承载远程直接内存访问（RoCE）网络技术。IB在市场中占据重要地位，但其产业开放性、部署成本面临挑战；而RoCE底层采用传统以太网，产业开放性好，但性能受到制约，基于传统以太网流式转发的优化，性能提升有限，且不适用于稀疏模型。随着AI大模型的快速发展，GPU集群规模从千卡向万卡、10万卡演进，稀疏模型也得到广泛应用。智算中心网络技术成为智算中心的瓶颈，需要从基本的转发机制进行革新，主要通过优化负载均衡和主动拥塞控制来实现网络高吞吐无损传输。

3) 跨园区智算中心间互联网络主要解决跨数据中心的服务器间通信问题。因基础设施条件受限，同园区资源无法支撑更大规模GPU卡的建设。借助已建智算中心的利旧方式，多个跨地域智算中心通过广域网互联共同支撑更大的集群训练。跨园区智算中心间互联网络主要承载跨智算中心服务器/机架间PP、DP数据，通过广域网互联，使用路由器与高速传输网络设备承载。因广域网存在带宽收敛、长距通信时延（5 μ s/km）以及安全隐患等问题，网络传输需通过优化负载均衡、拥塞控制及数据加密技术，实现安全、高吞吐的传输效果。同时，训练平台应采用分层调度优化等方式，减少跨数据中心的通信数据量，并通过计算尽量降低通信时延，将长距离传输的负面影响降至最低。目前，业界正积极探讨长距跨智算中心训练的可行性。

2 服务器/超节点内GPU互联网络

2.1 服务器/超节点内GPU互联网络的需求与挑战

随着人工智能技术的迅猛发展，超万亿参数的大模型逐渐成为颠覆型技术，在处理诸如语言理解、视频生成、科学计算等复杂任务时展现出超凡能力。然而，这种能力的发挥对计算架构提出了前所未有的挑战。万亿大模型的典型架构

是MoE结构。MoE通过将大模型分解为多个“专家”子网络，并按需动态调用这些专家，来达到以少量的算力资源获得更高模型能力的效果。尽管MoE在理论上能够显著提升模型的参数量并扩大计算规模，但在实际部署时仍依赖底层硬件提供强大的“专家”交互通道。当前，智算服务器只支持8卡直连拓扑（如图2所示），无法满足MoE模型高频率、大带宽、低延时的通信需求。为此，产业界正在积极开展超节点服务器的研究。这类服务器的核心，在于其能够利用服务器内的高带宽域（HBD），构建GPU间高效的Scale up互联，有效提升张量并行、MoE并行和流水线并行的传输效率，在不改变GPU芯片物理条件的前提下实现整体性能的优化。GPU Scale up互联的通信效率，尤其是带宽和延迟的优化，已成为业界的热点方向。

2.2 服务器/超节点内GPU互联网络技术业界发展情况

目前，GPU互联协议主要包括PCIe、Nvidia NVLink、AMD Infinity Fabric与UALink等。其中，PCIe由Intel公司于2001年提出，主要用于连接中央处理器（CPU）与硬盘、网卡等设备，是当前应用最广泛的外设互联协议。2022年1月，PCIe 6.0正式发布，其单通道速率达到64 GT/s，x16通道下双向传输速率达到256 GT/s。PCIe 6.0放弃了前代PCIe使用的不归零编码（NRZ）调制方式，采用四电平脉冲幅度调制（PAM4），为应对PAM4高误码率挑战引入了前向纠错（FEC）机制和循环冗余校验（CRC），以进一步提升数据传输的可靠性。PCIe可以支持多个GPU通过CPU或者PCIe Switch实现连接，但这种设计并不能满足日益提升的GPU互联需求：一是，GPU间无法点对点（P2P）通信，时延较长；二是，由于CPU提供的PCIe通道数量有限，当网卡等设备对PCIe通道的需求增加时，将挤占GPU占用的通道数，导致互联的GPU数量受限。此外，由于PCIe需要兼容多种低速设备，其带宽提升速度无法满足GPU通信需求，因此业界很少依赖PCIe搭建GPU间的高速互联通道。

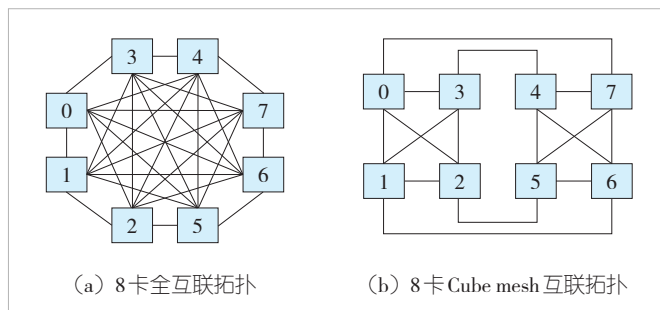


图2 传统8卡直连拓扑结构

支持GPU直接通信的互联协议可分为私有协议和开放协议两种。私有协议包括Nvidia NVLink和AMD Infinity Fabric。其中, NVLink是Nvidia在2016年P100产品中首次引入的。NVLink 1.0由成对的Sub-Link构成, 每个Sub-Link包含8个差分信号对, 以20 Gbit/s的速度传输NRZ的差分电信号, 实现单链路40 Gbit/s的双向带宽。P100支持4条链路, 总双向带宽达到160 Gbit/s。NVLink 1.0同时支持CPU-GPU以及GPU-GPU的P2P通信, 允许GPU直接读写远端CPU的主机内存和相连GPU的设备内存, 为高性能计算和AI应用提供了更强的数据传输能力。2024年3月, NVLink演进到了第5代, 并用于Nvidia GB200 NVL72超节点服务器。从架构上看, NVLink 5.0通过NVSwitch交换芯片实现了72个Blackwell GPU和36个Grace CPU的全互联, 如图3所示, 其中每个GPU的双向吞吐量达到1.8 TB/s, 是PCIe 5.0带宽的14倍以上。NVLink和NVSwitch为GPU集群带来了更高的通信带宽和更低的延迟, 从而提升了系统的整体性能和效率。

AMD的Infinity Fabric在设计之初主要用于实现CPU与CPU的高速互联。Infinity Fabric支持每条链路42 Gbit/s的双向带宽, 融合了数据传输与控制功能, 由可扩展数据网络(SDF)和可扩展控制网络(SCF)两个独立通信平面构成。其中, SDF负责计算核心、内存和I/O间的数据传输一致性, SCF提供系统配置和管理的通用命令与控制机制。2021年, Infinity Fabric升级到3.0版本, 并实现了CPU与GPU间的一致性互联, 支持CPU的DRAM内存与GPU的HBM内存的一致性内存架构。此外, 不同于英伟达NVLink仅限于内部使用, AMD已经逐步向其合作伙伴开放Infinity Fabric系统, 以完善技术生态布局。

2024年由AMD牵头主导, 亚马逊、Astera Labs、思科、谷歌、惠普、Intel、Meta和微软等行业巨头参与成立UAL-

ink联盟, 旨在开发一种新的开放行业标准, 专注于开发GPU卡间互联系统, 打破英伟达垄断。UALink初版协议来自AMD, 最多支持1 024个GPU互联。此外, UALink定义了创新的I/O架构, 具备高性能内存语义访问原生支持, 可实现显存共享、支持Switch组网模式、具备超高带宽和超低时延能力, 在性能和GPU互连规模上比肩Nvidia NVLink。UALink 1.0正式协议在2025年4月发布。

2.3 OISA核心技术与创新点

相比于Nvidia基于NvLink与NvSwitch提供256张GPU的高速互联能力, 国产相关技术目前仅支持8卡直连拓扑, 交换技术的缺失导致GPU无法向超节点演进。针对当前国产GPU卡间互联带宽低、时延高、互连规模受限等问题, 中国移动原创提出全向智感互联OISA协议体系, 研究GPU间的Scale Up互联技术, 定义物理层、数据层、事务层标准, 实现超节点内多GPU芯片的对等全互联。

OISA具有统一报文格式、多语义融合、多层次流控重传以及支持集合通信加速等关键技术特征, OISA协议栈及各层功能如图4所示。GPU片上互联可基于消息语义或内存语义实现。OISA同时设计了两套报文格式, 能够根据具体的应用场景选择最合适的通信模式, 从而优化性能和资源利用率。在报文可靠性传输方面, OISA引入事务层选择性重传和数据层重传两种机制。其中, 事务层选择性重传功能基于滑动窗口实现, 能够精确识别并仅重传在互联过程中丢失或损坏的数据包, 而不是无差别地重传整个数据流。与传统Go Back to N重传机制相比, OISA通过选择性重传机制提高了传输效率并减少了资源浪费, 有助于实现高效物理传输设计和系统扩展性的目标。数据层重传技术能够在数据层检测到错误, 并立即触发表点对点的重传机制, 避免错误影响整个协议栈, 从而实现数据传输的低时延和高可靠性。在流量控制方面, OISA使用了流量感知机制, 支持基于信用的流量控制(CBFC)和

基于优先级的流量控制(PFC)两种流量控制机制, 通过实时监控和流量模式分析, 动态调整资源分配和数据传输速率, 确保数据流的最优分配和信息提取。CBFC和PFC分别基于信用额度和数据优先级来精确控制数据流, 有效减少网络拥塞和丢包的可能性, 提高数据传输的效率和可靠性。最后, OISA支持集合通信加速技术, 并将部分集合通信加速能力卸载至交换芯片侧, 减少GPU与Switch的频繁数据搬运, 降低模型训练和推理的时延。

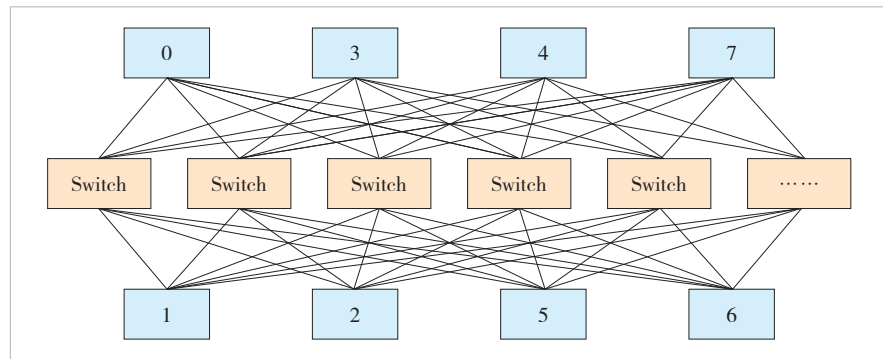


图3 基于Switch的交换拓扑结构

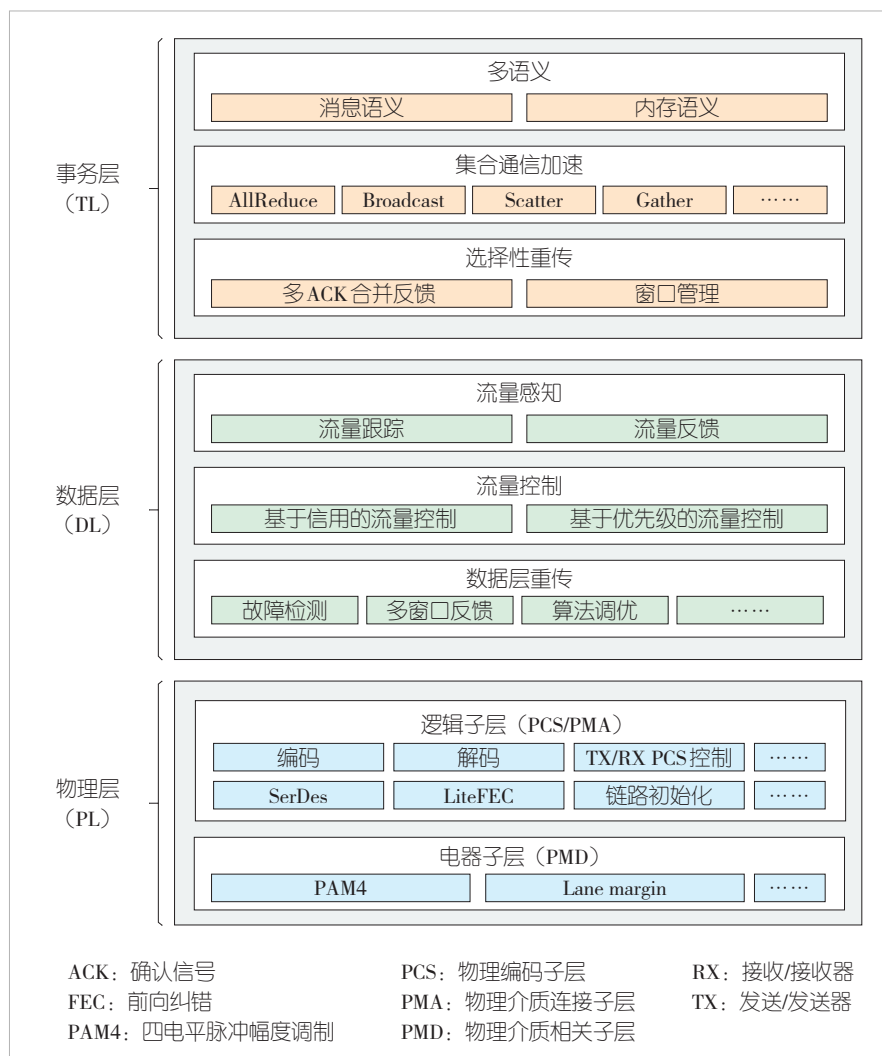


图4 OISA协议栈及各层功能

3 同园区机间互联网络

3.1 同园区机间互联网络需求与挑战

智算中心机间网络主要用于承载AI模型训练业务的PP和DP交互的权重、梯度等数据，GPU算力性能高，接入网络带宽100 Gbit/s以上，并通过RDMA来降低传输时延。因此，大模型训练的智算中心网络需要缩短迭代过程中通信传输数据的时间，降低通信开销，从而减少GPU的计算等待，提升计算效率。同时，机间GPU通信流量呈现周期性、同步突发的特点。在大模型训练过程中，通信具有非常强的周期性，且每轮迭代的通信模式均保持一致。在每一轮的迭代过程中，不同节点间的流量保持同步，同时流量以On-Off的模式进行突发式传输。一旦网络出现拥塞或故障，可能会导致数据丢失或传输延迟，从而影响AI模型的训练时效。以上通信流量的特点对机间智算中心网络提出了高带宽、高

吞吐、低延时、高可靠的要求。

同园区智算中心GPU之间通过交换机CLOS组网架构实现互联，每个GPU接入网络带宽需求在100~400 Gbit/s之间，且同步并发要求CLOS组网1:1不收敛。随着GPU性能的增强和GPU集群规模的增加，同园区智算中心GPU对网络的单端口带宽、节点间的可用链路数量及网络总带宽提出了很高的需求。

基于RoCE协议的智算中心网络，通常采用五元组哈希实现链路负载分担技术，并通过PFC、显式拥塞通告(ECN)协议实现网络无损。然而，这种方案在智算中心网络应用中面临4个挑战^[5]：

挑战1：传统基于逐流的等价多路径路由(ECMP)负载均衡技术，在智算中心网络单流量规模大、流量数量小的情况下会失效，导致流量在交换网络发生极化。这不仅造成链路负载失衡，还会导致部分物理链路因承载过量流量而产生拥塞，最终引发报文丢弃问题，严重影响网络传输效率与稳定性。

挑战2：当前流量进入网络时，在不考虑出端口转发能力的情况下，流量会以“推”的方式进入网络。分布式训练的多对一通信模型会产生大量的Incast流量，造成设备内部队列缓存瞬时突发，进而引发拥塞甚至丢包。PFC和ECN都是

拥塞产生后的被动拥塞控制机制，它们无法从根本上避免拥塞，且1个比特的ECN信号仅能定性地表示网络产生拥塞，无法定量地表示拥塞程度。对此，端侧需要探测式地调整发送速率，但收敛速度慢，容易导致网络吞吐性能下降。

挑战3：由于网络故障无法避免，因此高效的故障管理和快速的业务恢复机制至关重要。智算中心应配备先进的故障预测工具与自动化故障恢复系统。这些系统可以基于历史数据和实时性能指标，预测并识别潜在的故障点。一旦检测到故障，自动化的故障解除和业务恢复流程就可以迅速启动。这样可以减少系统的停机时间，确保业务的连续性。

挑战4：当前AI超大模型的参数已达千亿甚至万亿级别，训练大模型需要超高算力，并且对显存需求也很高，而更高的显存消耗意味着需要千卡/万卡GPU才能完整存储一个模型的训练过程。组网规模的大幅增长导致网络管理更加

复杂, 拥塞控制、负载均衡的难度也会增加。

3.2 同园区机间网络技术业界发展情况

传统智算中心网络技术主要包括 InfiniBand 和 RoCE 两种技术路线。其中, InfiniBand 使用专有的硬件和优化协议, 产业生态链成熟, 但是生态封闭, 网络的建设和维护成本相对较高; RoCE 基于标准以太网, 可以利用现有的以太网基础设施, 降低了网络建设的成本, 但性能受到制约。随着 AI 大模型的快速发展, 智算中心网络技术已经成为全球人工智能巨头关注的焦点, 其核心是新一代以太网技术的突破。

2023 年 5 月, 中国移动携 50 余家业界伙伴发布了全调度以太网技术 (GSE)^[6], 原创提出新型以太网转发和调度机制, 以国产技术和产业生态为主, 构造低时延、零丢包的无损网络。当前智算中心服务器主要有两类: 一类是 GPU 集成网卡, 典型产品如华为昇腾 910 系列; 另一类是配备独立网卡的 GPU 服务器, 典型产品如英伟达 H800 系列, 通常需要不同的网络解决方案。GSE 采用统一设计理念, 形成了 GSE-N2N 和 GSE-E2E 两大模式, 以满足各种智算中心的网络需求。其中, GSE-N2N 技术方案适用于 GPU 集成网卡场景, 将网络设备作为全调度网络处理节点 (GSP) 和全调度交换网络 (GSF) 节点, 支持 GSE 的全部功能, 使得 GPU 服务器与 GSE 网络之间无须直接联动, 就能实现天然解耦, 同时确保无损、高性能的集群互联; GSE-E2E 技术方案则适用于配备独立网卡的 GPU 服务器, 通过将部分 GSE 能力延伸至 GPU 服务器的网卡, 将网卡作为 GSP 节点、网络设备作为 GSF 节点, 借助端网协同实现高性能的集群互联。

2023 年 7 月, 博通、思科、Arista、微软、Meta 等美国巨头企业牵头成立超以太网联盟 (UEC)。该联盟秉持生态开放的理念, 将技术路线重心放在以太网通信协议栈的升级上, 通过优化物理层、链路层、传输层和软件应用程序编程接口 (API) 层等方面的技术, 提升以太网在高性能计算和网络通信方面的能力, 其目标是打造面向 AI 时代的超大规模新型网络系统。UEC 的目标愿景和技术理念与中国移动的 GSE 一致。

3.3 GSE 核心技术与创新点

GSE 革新以太网底层转发机制和上层协议栈, 创新提出报文容器 (PKTC) 的多路径喷洒、动态全局调度队列 (DGSQ) 的主动拥塞避免等关键技术, 突破无损以太性能瓶颈。主要创新点包括:

1) 基于报文容器的负载均衡机制提升网络吞吐

在智算中心流数少、单流量大、多流并发的场景下, 传统基于流负载分担的方式存在网络拥塞和吞吐率低的问题。GSE 通过在源 GSP 节点将数据包逻辑组装成固定长度的报文容器, 并为每个报文打上 GSE 头、容器标识等。源 GSP 和中间 GSF 以报文容器粒度进行多路径负载均衡, 目的 GSP 节点以报文容器为粒度进行重排, 实现报文保序交付, 使网络吞吐率高达 99%, 如图 5 所示。

2) 基于 DGSQ 的请求授权机制避免拥塞

智算中心存在多打一 Incast 场景, 传统被动拥塞控制有一定的滞后性, 难以有效避免丢包和时延增加。GSE 通过动态调度授权机制避免拥塞, 调度授权流程如图 6 所示, 在源 GSP 接收报文后, 向该报文对应的目的 GSP 进行信令申请; 目的 GSP 收到各个源 GSP 的信令申请后, 按照周期或按次进行信令分配, 在极致负载情况下, 主动控制进入网络的流

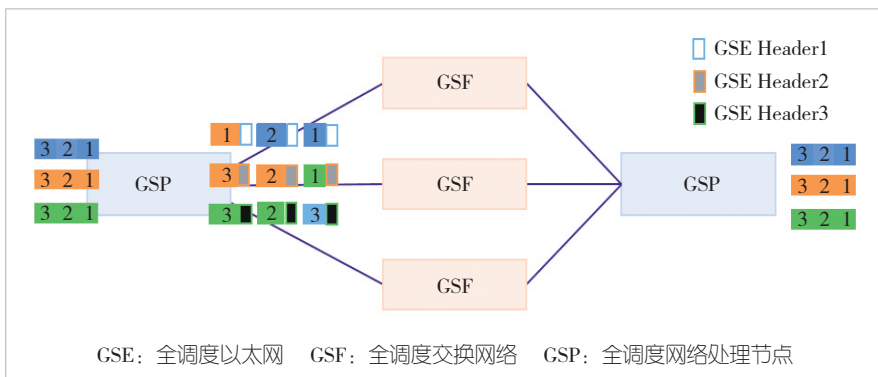


图5 报文容器转发示意图

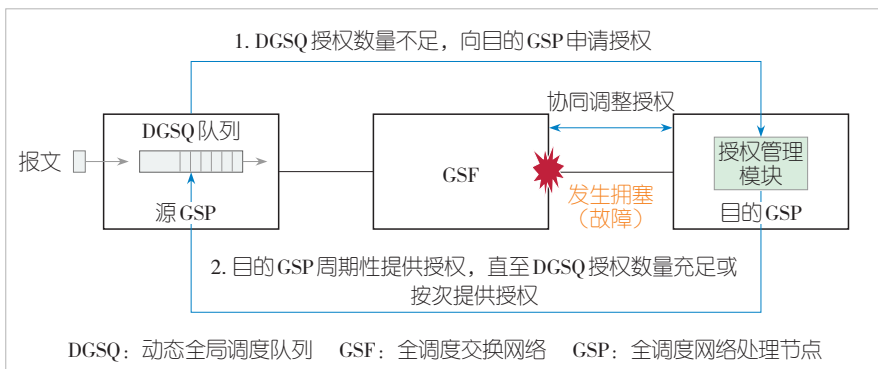


图6 DGSQ调度授权流程

量，实现网络原生无损，并根据实际流量特征按需创建授权队列资源，支持大规模扩展。

3) 全局+分布式控制面实现网络设备即插即用

全调度以太网控制面分工如图7所示。集中式全调度操作系统（GSOS）用于全局信息编址（例如设备节点ID等）、设备IP地址配置等。设备端网络侧操作系统（NOS）具备独立的控制面和管理面。GSP/GSF通过eBGP路由打通网络连接，并通过多协议边界网关协议（MP-BGP）在设备间交互控制信息。各设备分布式构建全局转发表，实现GSE报文转发、运行容器负载均衡、DGSQ调度等GSE网络功能，并通过控制器集中式管理和设备分布式控制，实现网络设备的即插即用。

4) 微秒级感知网络故障并收敛保证无损

基于报文容器的负载均衡与DGSQ的请求授权机制在正常网络环境下可实现零丢包转发，但当网络发生链路中断或设备故障时，拓扑结构的非对称化仍可能引发局部拥塞。为此，GSF通过光电信号通断检测与误码率实时监测实现微秒级故障感知，并基于66B原子码块或以太报文进行全局通告，将故障通告报文传递至所有关联的GSP。各GSP收到故障报文后进行负载均衡收敛并调整授权，从而保证网络在故障场景下尽力无损，如图8所示。

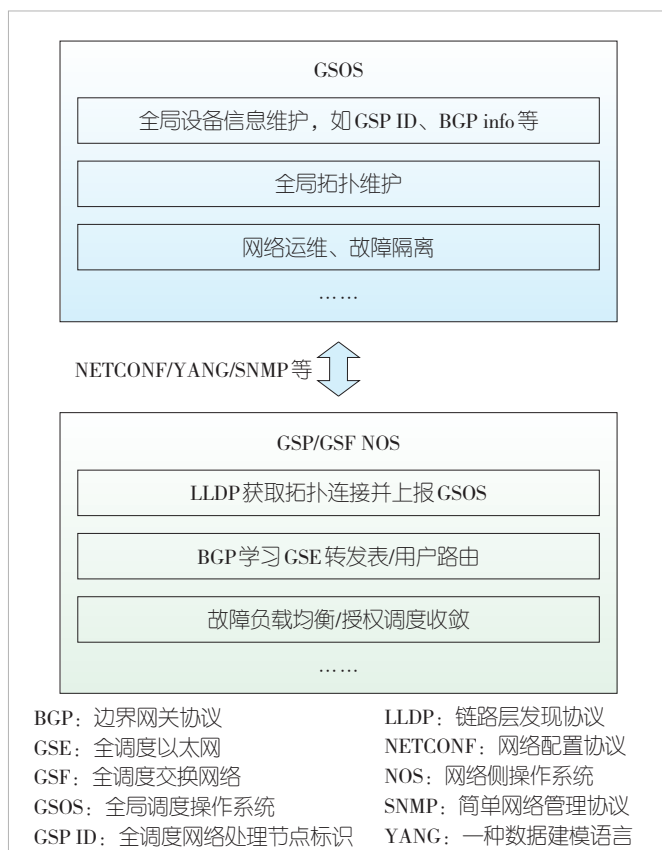


图7 全调度以太网控制面分工

3.4 同园区机间网络组网建议

根据园区组网规模和机房条件，园区内单集群建设分为同机房、同楼宇跨机房、同园区跨机房3种场景。为减少对GPU集群算力效率的影响，并简化工程部署设计和运行管理的复杂度，所有GPU服务器都尽量部署在同一个机房内。当基础设施条件不能满足工程部署要求时，可把一个智算集

群分散到多个机房进行跨机房部署，按照机房容量从大到小的原则，充分利用单个机房的容量资源，尽量减少跨机房部署的服务器和网络数量，缩短拉远的布线距离。当在同一个楼宇内找不到满足工程部署条件的机房时，需要考虑采用同园区跨楼宇的拉远部署方式——在园区内的不同楼宇间预埋

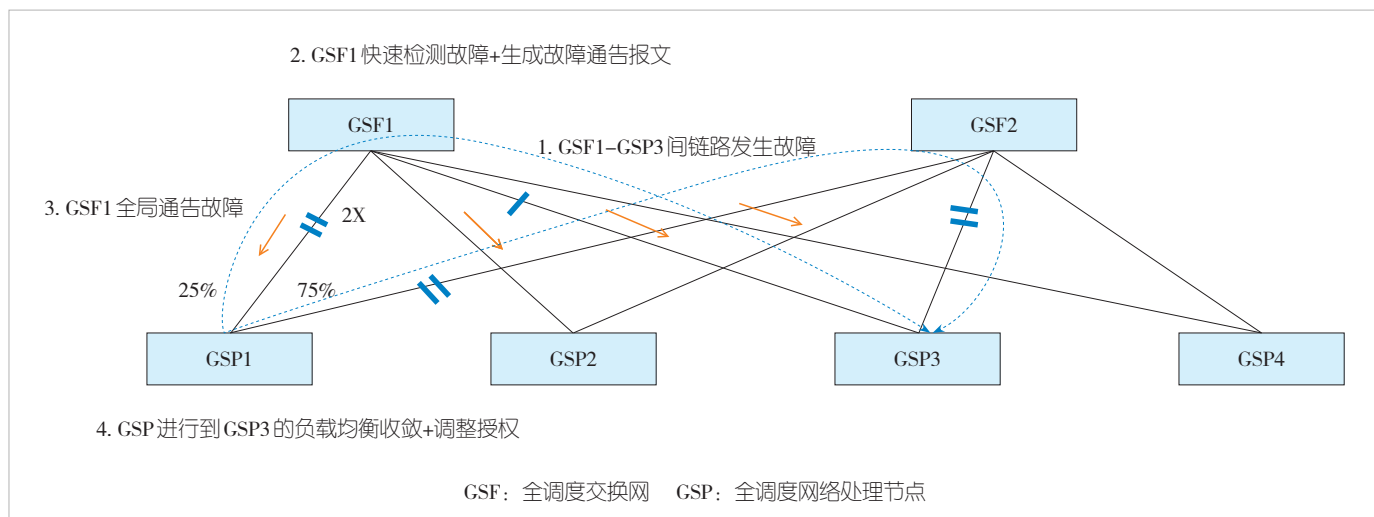


图8 网络可靠性保证流程图

好相应数量的裸纤，以便进行拉远部署。

考虑到交换机转发芯片容量和建设节奏，建议按照POD为单位建设，每个POD接入数千卡至万卡规模，通过平滑扩展POD建设超万卡/十万卡规模的集群。POD间通过高速交换网互联，POD内和POD间根据业务需求可设置一定的收敛比。PP通信量小，但是对时延敏感，在带宽收敛比大、静态时延小的情况下优先跨POD。DP对时延不敏感，但是并行通信量大，在带宽收敛比小、静态时延高的情况下优先跨POD。DP跨POD时可借助训练框架和集合通信算法感知POD，通过分层AllReduce算法进行优化，以减少跨POD间的通信量。

4 跨园区智算中心互联网络

4.1 跨园区智算中心互联网络需求与挑战

AI集群算力需求的爆发式增长，催生了以网补算的协同训练需求。为突破单中心算力瓶颈，业界开始探索分布式智算中心协同训练模式，将训练任务分散至多个中心并行计算，以扩大整体算力规模。同时，为提升智算中心资源利用效率，服务商期望整合多中心资源，实现AI算力的集中管理与资源池化。然而，跨智算中心互联打破了单体集群的诸多硬件和组网设计前提，导致带宽、时延等网络指标恶化，给网络带来了极致性能挑战。

此外，海量AI训练、推理样本的产生，使得数据网络传输需求呈指数级增长。以乘用车自动驾驶训练为例，每辆路测车每天产生5~10 TB的样本数据，且通常要求在24 h内从分布于不同城市的路测基地传输至智驾训练中心。按照500辆路测车计算，其路测样本数据量每年超过600 PB，约相当于6亿部高清电影的数据量。这对网络带宽、延迟和吞吐量提出了更高要求。因此网络必须具备弹性高带宽能力，以应对大规模数据传输的挑战。

智算中心以大量数据为资源，涉及大量数据在入算、算内和算间网络场景的处理和传输。这些数据是企业十分重要的商业资产，一旦被窃听攻击或泄露，将产生难以估计的经济损失。在跨园区智算中心互联场景下，用于传输智间数据资源的高速互联光纤链路以及相关设施暴露在物理环境中，使得数据窃听、攻击风险倍增，网络面临严峻安全挑战。

4.2 跨园区智算中心网络业界技术发展

随着智算中心规模的不断扩大以及算力需求的指数级增长，如何实现高效、稳定、低成本的智算互联，成为了行业发展的关键问题。

思科在跨园区智算中心互联领域提出了路由光网络技术方​​案，通过利用标准化的400G ZR/ZR+可插拔光模块，直接从路由器端口输出相干波长信号，实现了IP层与光层的深度融合。这种架构不仅能简化网络设计，降低运营成本，还可满足智算互联场景对高性能和高可靠性的需求。谷歌积极布局多数据中心分布式训练，其多数据中心间距离最近约为24 km，最远约为80 km。预计到2026年，Google将形成4个园区的跨DC智算互联集群。其Gemini1Ultra系列模型通过光交换机实现多数据中心间的高效、低成本互联^[7-9]，网络延迟和带宽足以支持常用的同步训练模式，在超级Pod内完成模型并行，在Pod间完成数据并行。诺基亚在智算中心互联领域推出了创新的IP与光融合解决方案，针对智算互联场景中多样化的数据传输距离和性能需求提供了灵活的组网方案：在120 km范围内的400G互联场景中，可通过路由器配置灰光或彩光模块实现高效连接；对于更长距离的传输需求，可采用路由器配置400G多跨光模块方案；而对于多段超长距离的传输，则可通过光网络系统的400G光传输技术实现无缝连接。

腾讯与中国移动正在开展IP+光融合的智能互联技术研究，旨在为智算中心提供高效、灵活的网络互联解决方案^[10]。在该方案中，通过采用高性能可插拔光模块，并将其直接部署在智算中心互联网络设备上，网络设备能够直接输出高性能彩光信号。这些信号可以直接进入光层系统的合波器，形成密集波分复用（DWDM）合波信号，无须依赖光网络中的OTU单元进行信号的复用和映射，从而显著简化了网络架构。此外，针对IP over DWDM技术，腾讯正在推动光器件的小型化与标准化，即微光学模块的发展，以进一步提升系统的灵活性和可扩展性。传统“路由器+光模块+光传送网+DWDM”的复杂网络架构被简化为“IP+光融合”网络。这种方案不仅降低了处理时延，还减少了多组电层设备的使用频次，使传输路径更加简洁、部署成本显著降低，同时大幅简化了系统部署流程，为智算互联场景提供了高效、低延迟、低成本的基础设施支持。

阿里云在智算中心互联领域通过自研的eCore网络架构，采用单栈单片设计（基于IPv6/SRv6协议栈和单芯片白盒路由器）和多平面架构，实现了全网服务化的技术突破^[11]。eCore深度融合IPv6/SRv6协议栈与白盒硬件，结合服务化设计，显著提高了网络资源利用率，降低了业务部署复杂度。该架构为智算互联场景提供了高带宽、低延迟的网络基础设施支持，能够高效满足大规模AI训练、分布式计算以及全球化业务的需求，助力智算中心实现高性能、高可靠的网络互联，同时推动智算互联技术的快速发展和规模化应用。

综上,面向智算互联场景,IP+光融合技术日益受到业界关注,并逐渐成为解决高带宽、低延迟网络需求的关键方向。业界在这一领域已开始逐步应用,通过将IP层与光层深度融合,显著提升了网络性能和效率。

4.3 跨园区智算中心互联创新关键技术

随着数据流量的爆炸式增长、新兴应用的涌现以及对网络性能要求的不断提升,传统网络架构面临巨大挑战。在此背景下,智算集群互联技术组合应运而生,为构建高效、智能、可靠的网络基础设施提供全新机遇。IP与光通信技术的深度融合,旨在充分发挥两者优势,实现高性能、高带宽、低延迟、高安全和高可靠的网络,以满足日益增长的数字化需求。面向智算互联场景,中国移动在IP广域网络技术上全面创新,革新现有流量分发、拥塞控制和安全加密机制,提出弹性以太网聚合FlexPipe、微流级拥塞控制以及物理层安全加密PHYSec等关键技术,助力构建大运力、无阻碍、高安全的跨智算中心互联网络。主要创新点包括:

1) 弹性以太网聚合FlexPipe技术

当前智算中心互联的路由器设备速率仅为100 Gbit/s,当GPU面临超大AI训练流时,数据中心互联路由器的100G端口必然会出现拥堵丢包的情况。传统广域网采用五元组Hash负载分担方式,常常会将大象流分配至单一链路,即便其他链路仍有空闲资源,该链路也会不可避免地出现拥塞状况,进而产生丢包现象。基于报文组调度的弹性以太网聚合FlexPipe技术能够依据用户的带宽需求,灵活地实现不同带宽速率的以太网端口无损转发。在不依赖特定硬件的情况下,单设备最大可达3.2T的超宽通道。此外,它还能在超大的FlexPipe管道内部,为多条物理端口间的大象流提供负载分担能力,进一步提升了网络对高带宽突发流量情况的带宽灵活适应能力,为智算中心间的流量无损转发提供坚实保障。

2) 微流级拥塞控制技术

在通过跨智算中心的分布式训练或存算分离等方式进行模型训练时,智算中心资源池间距离的拉远会导致传输时延增加,从而影响网络状态反馈的及时性。然而,智算业务具有超大流、同步传输、任务式、丢包敏感等特征,在进入广域网时,对网络传输的可靠性提出了极高要求。特别是在丢包问题上,智算业务表现出极强的敏感性。因此,需要构建具备稳定连接能力的网络基础设施,以确保智算业务的高效运行。微流级拥塞控制技术能够根据实时网络状况动态调整流控策略,实现流量峰值速率流级别的独立控制和精准反压。相较于传统的PFC机制,微流级拥塞控制技术解决了传

统PFC的头阻、反压风暴和死锁问题,同时基于精准的IP大象流识别和预测,实现了对网络中每一条流的独立监控与动态调整,通过以数据流为单位的精准流量控制并结合大缓存设计及调度,可实现长距网络环境下的业务“零”丢包,将拥塞和故障带来的影响最小化。该技术能够有效解决由短暂拥塞导致的潜在问题,确保数据传输的连续性和完整性。

3) 物理层安全加密PHYSec技术

现有的媒体访问控制安全(MACSec)加密机制是基于数据链路层的安全加密技术,可以保证以太网设备之间数据帧的安全性。但加密后的数据帧仍会暴露以太帧头部信息,同时不能对优先级流量控制帧或Pause帧进行加密。此外,MACSec加密认证过程也会引入较大的封装开销,占用一定的业务带宽。对此,中国移动原创提出的全球首个以太网物理层安全PHYSec技术,通过将传统密码学的理念与以太网物理层技术相融合,提供链路级加解密能力与通道级加解密能力。如图9所示,链路级加解密技术将多个对齐标识(AM)数据段复合形成复帧比特流,作为PHYSec的最小加解密单元,优先在模块内部署。通道级加解密技术针对64B/66B码块流进行全部加密,增加D码块承载解密所需的参数,需要在物理层芯片内部署。该技术具备低时延、高安全、低开销和协议透明等优势,满足数据链路层及所有上层协议的信息防护要求。

5 结束语

随着生成式人工智能的蓬勃发展,新型智算网络技术已成为全球创新焦点。智算中心网络是一个多要素融合的复杂系统,是算网的深度融合。中国移动携业界伙伴创新提出OISA、GSE技术体系,以及弹性以太网聚合、精细化拥塞控制、物理层安全等多项创新技术,旨在构建一个标准开放的高性能智算网络技术体系,助力人工智能生成内容(AIGC)业务快速发展。

中国移动已发布《全向智感OISA核心规范》Gen1.1版本协议,GSE全套技术标准《GSE1.0算网协同技术标准》《GSE2.0端网协同技术标准》《GSE2.0网络侧优化技术标准》,以及《新型智算中心以太网物理层安全(PHYSec)架构白皮书》。OISA可支持128张GPU互联,点对点带宽达到800 Gbit/s,极大提升了国产GPU的互联能力。GSE试点表明,OISA的组网性能比传统RoCEv2交换机提升50%以上,达到国际先进水平。

未来,中国移动会持续丰富OISA、GSE协议体系,加速跨园区互联技术创新,并推动构建产业化生态。技术标准的落地依赖AI业务、机内/外交换芯片、网卡、网络设备等

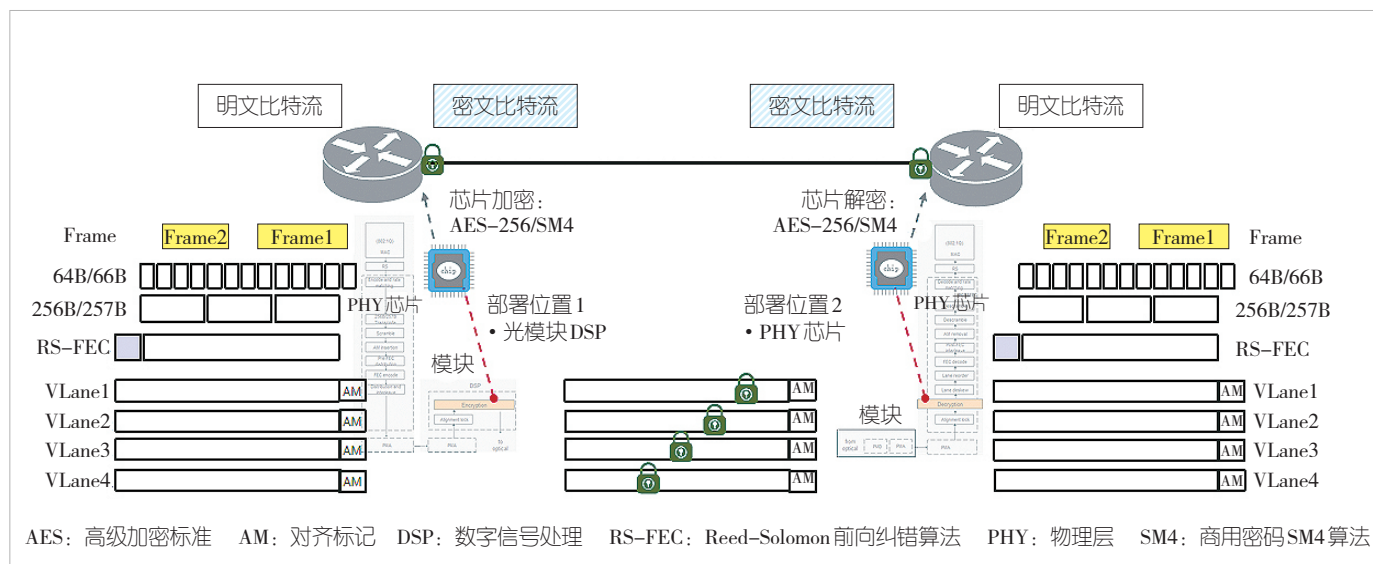


图9 PHYSec技术架构

上下游的协同创新。由于创新难度大、开发周期长，我们希望产业各界能够携手合作，协同创新，共同推动技术进步，为产业发展贡献力量。

算力网络新项目 [EB/OL]. (2024-12-17) [2025-03-10]. https://mp.weixin.qq.com/s/k_7xtw-1uMUbK7_Km1hcrw
[11] 阿里云基础设施. 阿里云基础设施网络 2024 年创新总结 [EB/OL]. (2025-01-20) [2025-03-10]. <https://mp.weixin.qq.com/s/JnziKrk8dYcOcafj29apA>

参考文献

- [1] KAPLAN J, MCCANDLISH S, HENIGHAN T, et al. Scaling laws for neural language models [EB/OL]. [2025-03-10]. <https://arxiv.org/abs/2001.08361v1>
- [2] WEI J, TAY Y, BOMMASANI R, et al. Emergent abilities of large language models [J]. Transactions on machine learning research, 2022(8): 1-30
- [3] DU Z X, ZENG A H, DONG Y X, et al. Understanding emergent abilities of language models from the loss perspective [EB/OL]. [2025-03-10]. <https://arxiv.org/abs/2403.15796v3>
- [4] HOWARTH J. Number of parameters in GPT-4 [R]. 2024
- [5] 中国移动通信研究院. 面向 AI 大模型的智算中心网络演进白皮书 [R]. 2023
- [6] 段晓东, 程伟强, 王瑞雪, 等. 面向新型智能计算中心的全调度以太网技术 [J]. 中兴通讯技术, 2023, 29(4): 57-63. DOI: 10.12142/ZTETJ.202304011
- [7] Gemini Team. Gemini: a family of highly capable multimodal models. Google [EB/OL]. [2025-03-10]. https://storage.googleapis.com/deepmind-media/gemini/Gemini_1.0.pdf
- [8] HONG C Y, MANDAL S, AL-FARES M, et al. B4 and after: managing hierarchy, partitioning, and asymmetry for availability and scale in Google's software-defined WAN [C]//Proceedings of the 2018 Conference of the ACM Special Interest Group on Data Communication. ACM, 2018: 1-44. DOI: 10.1145/3230543.3230545
- [9] POUTIEVSKI L, MASHAYEKHI O, ONG J, et al. Jupiter evolving: transforming Google's datacenter network via optical circuit switches and software-defined networking [C]//Proceedings of the ACM SIGCOMM 2022 Conference. ACM, 2022: 66-85. DOI: 10.1145/3544216.3544265
- [10] CAICT 算力. MegaScaleOut: 聚焦百万量级 GPU 集群, ODC 启动

作者简介



段晓东, 中国移动通信有限公司研究院副院长, 教授级高工, 享受国务院特殊津贴专家, 工信部通信科技委信息通信专家组专家, “新世纪百千万人才工程”国家级人选; 长期从事 AI 智算网络、下一代互联网、5G 网络架构、SDN/NFV 下一代网络及 6G 网络研究等工作, 先后主持了多项国家重大专项课题, 并获得国家科技进步奖特等奖等多个奖项; 获授权专利 50 余项。



程伟强, 中国移动通信有限公司研究院基础网络技术研究所副所长, 教授级高工; 长期从事下一代互联网、数据中心网络、传输网等方面的技术研究和标准推动工作。



张昊, 中国移动首席专家、中国移动通信有限公司研究院网络与 IT 技术研究所所长, 高级工程师; 长期从事 4G/5G 核心网、边缘计算、算力网络、新型智算等 ICT 融合领域技术研发工作。