



专栏

智算中心存储系统优化研究

曹原铭, 雷鸣, 刘芹, 牛瑛霞, 武振宇, 潘洁
(中国移动通信集团设计院有限公司, 北京 100080)

摘要: 智算中心使用分布式文件存储进行数据预处理和模型训练, 使用分布式对象存储进行原始数据获取和模型发布, 使用分布式块存储为资源管理平台提供存储服务。同时, 模型训练过程中经常使用高性能分布式文件存储缩短 checkpoint 读写时间, 提高集群训练效率。大模型训练全生命周期需要在不同存储协议、不同读写性能的存储系统之间进行数据复制、迁移, 导致数据重复存储, 且数据复制需要占用计算资源和网络带宽。为了解决上述问题, 并为智算集群提供统一命名空间, 提出文件、对象融合存储和文件系统分级存储方案, 解决不同存储协议间的数据搬运问题, 实现高性能文件存储(全闪)和普通性能文件存储(混闪)间的数据自动流动, 为超大规模智算集群的存储系统优化方案提供参考。

关键词: 融合存储; 分级存储; 分布式文件存储; 分布式对象存储

中图分类号: TP393

文献标志码: A

doi: 10.11959/j.issn.1000-0801.2025160

Research on optimization of storage system in intelligent computing center

CAO Yuanming, LEI Ming, LIU Qin, NIU Yingxia, WU Zhenyu, PAN Jie
China Mobile Communications Group Design Institute Co., Ltd., Beijing 100080, China

Abstract: The intelligent computing center uses distributed file storage for data preprocessing and model training, distributed object storage for the acquisition of raw data and model release, and distributed block storage to provide storage for the resource management platform. Meanwhile, high-performance distributed file storage is used to shorten the read and write time of checkpoint during the training process and improve the training efficiency of the cluster. The entire life cycle of large model training requires data copying and migration between storage systems with different storage protocols and different read-write performances, resulting in duplicate data storage. Additionally, data copying requires computing resources and network bandwidth. To address the above issues and provide a unified namespace for the intelligent computing clusters, a file and object converged storage and a hierarchical file storage scheme were proposed to solve the problem of data transfer between different storage protocols and enable automatic data flow between high-performance file storage (all SSD) and ordinary-performance file storage (SSD and HDD), providing a reference for the optimization of storage systems of ultra-large-scale intelligent computing clusters.

Key words: converged storage, hierarchical storage, distributed file storage, distributed object storage

收稿日期: 2025-03-21; 修回日期: 2025-06-27

通信作者: 潘洁, panjie@cmdi.chinamobile.com

0 引言

智算中心采用人工智能 (artificial intelligence, AI) 服务器为经济和社会发展提供智能算力服务, 有力促进了 AI 产业化、产业 AI 化及政府治理智能化。当前各国政府都在加快推进 AI 发展, 加速从“互联网时代”迭代升级为“人工智能时代”, 我国政府提出深化大数据、AI 等研发应用, 开展“AI+”行动, 打造具有国际竞争力的数字产业集群, 这使得智能算力需求急剧增加。

业界主流大模型快速迭代演进, 参数规模由千亿级进入万亿级, 模型由单模态进化为多模态^[1], 面向终端消费者 (toC) 类高价值应用快速落地, 面向企业 (toB) 类行业应用逐渐起步, 深入赋能全社会各个领域, 促使全社会进一步加快“AI+”转型发展。随着模型参数量急剧增加、训练数据形态丰富多样、序列长度持续增大, 大模型训练对算力和存储的资源规模提出了更高的需求。DeepSeek 大幅降低大模型的使用门槛, 加速 AI 应用的普及, 推动行业的进一步繁荣, 带动算力、存力需求的持续增长。国内外科技公司更是加大智算中心资源布局, 行业头部企业正在加快构建万卡、超万卡集群的智算中心。

1 智算中心存储系统现状

1.1 智算中心存储系统方案

智算中心通常需要同时配置分布式对象存储^[2-3]和分布式文件存储^[4], 以满足 AI 训练数据获取、数据预处理、模型训练和模型发布等全生命周期的数据读写需求^[5], 同时为了缩短训练过程中 checkpoint 读写时间, 减少智算服务器等待时间, 提高集群训练效率, 通常会配置全闪文件存储。现有智算中心存储系统部署示意图如图 1 所示。

图 1 中共涉及对象存储、混闪文件存储和全

闪文件存储 3 种类型, 3 种存储分别独立配置软硬件资源, 不同类型存储资源之间的数据不会自动流动。对象存储和混闪文件存储采用 25 GE 端口搭建普通以太网, 为智算集群提供数据读写服务; 全闪文件存储采用 100 GE 端口搭建基于融合以太网的远程直接内存访问 (remote direct memory access over converged Ethernet, RoCE) 网络协议^[6], 为训练过程中的 Checkpoint 等提供高性能数据读写服务。

对象存储、文件存储等多种存储协议分别部署存储池, 以及不同性能 (全闪、混闪) 的分布式文件存储独立部署的方式, 可以基本满足千卡及以下规模的智算集群在训练过程中的数据读写性能需求, 但这种方式也存在一些问题, 如不同存储池间的数据重复存储、搬迁和复制等, 这些问题在万卡及以上的超大规模智算集群中更为突出。

1.2 存储系统存在的问题

1.2.1 资源配置效率低, 数据搬迁耗时长

AI 训练的不同阶段需要使用不同的非结构化数据存储协议, 如简单存储服务 (simple storage service, S3)、网络文件系统 (network file system, NFS)、Hadoop 分布式文件系统 (Hadoop distributed file system, HDFS^[7])、可移植操作系统接口 (portable operating system interface of UNIX, POSIX^[8]) 等协议。当不同协议的存储池分别配置时, 需要分别进行规划和建设, 无法直接复用一套存储系统来满足全流程业务场景。同时, 在训练过程中, 数据需要在不同协议的存储池之间频繁复制, 这不仅会占用计算、网络带宽等资源, 浪费存储空间, 还会造成大模型训练过程中智算服务器等待, 进而增加训练时长, 影响训练数据处理速度和模型训练效率。分立存储池时数据复制示意图如图 2 所示。

1.2.2 冷热数据无法自动流动

智算集群通常同时配置混闪文件存储和全闪

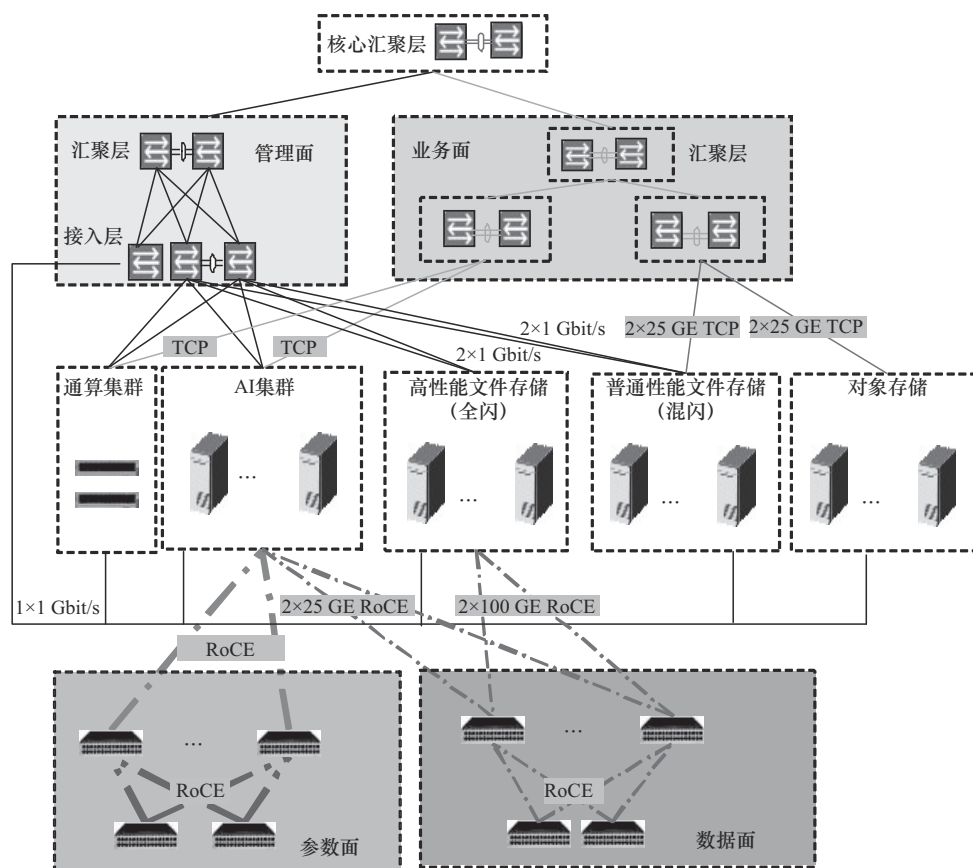


图1 现有智算中心存储系统部署示意图

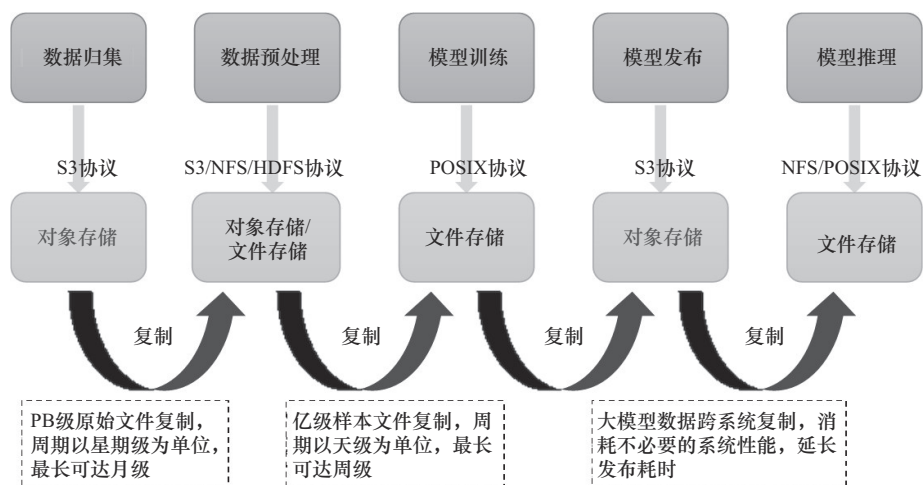


图2 分立存储池时数据复制示意图

文件存储，兼顾存储容量和存储性能，平衡训练效率和成本效益。两类文件存储独立配置时，冷热数据无法根据预置策略在两个存储池间自动流动，需要用户根据不同训练阶段的数据访问情况规划数据存储位置，维护数据一致性。同时，两个存

储池间的数据复制需要占用通算服务器和智算服务器的计算资源，峰值影响超过10%。在当前的部署方案中，数据在全闪和混闪之间流动会占用计算网络和存储网络带宽资源，影响其他业务的处理效率。

2 存储系统优化方案

针对不同存储池分立导致的数据跨池复制耗时长、冷热数据无法自动流动等问题,本文提出融合存储和分级存储方案,旨在提升智算中心存储系统的服务效率,为大模型训练加速。

2.1 融合存储

为了解决不同存储协议分别配置存储池带来的数据频繁复制问题,以及这一问题对训练效率的影响,可以采用文件与对象协议融合的存储方案,一个存储池同时具备文件和对象协议访问能力,不需要再单独规划文件和对象存储池,不仅能节省存储容量,还能实现训练全流程数据零复制,从而提升数据使用效率和训练效率^[9-10]。

2.1.1 融合存储的当前实现方式

业界现有的文件和对象融合存储产品多以文件或对象存储为底座,辅以协议转换网关方式实现。融合存储技术(通过网关实现)原理示意图如图3所示。

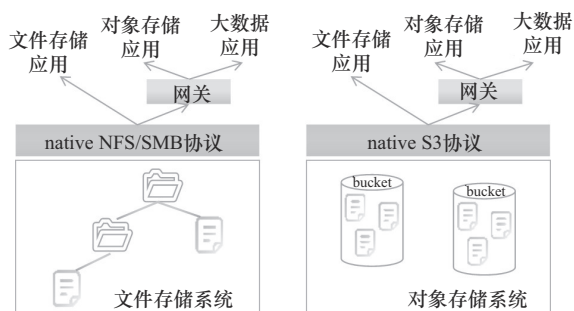


图3 融合存储技术(通过网关实现)原理示意图

融合存储技术通过在存储系统上层增加协议转换网关来解决不同协议互通的问题,可以同时提供文件和对象存储服务,避免跨协议数据复制。然而,网关设备增加了成本,协议转换影响了存储系统的性能,限制了部分功能,如以文件存储系统为底座时,存在对象协议难以支持多版本、多段上传、字典序等功能,单桶中支持对象数受限于单目录文件数能力,对象列举操作性能

差等问题;以对象存储系统为底座时,存在不支持目录权限修改操作、不支持文件锁、无法保障跨客户端操作的数据一致性等问题。

其他融合存储产品或技术同样无法实现存储设备共享,或需要对计算侧进行改造。文献[11]中提到的超融合存储产品底层为文件、块和对象配置不同的硬盘,组成3个相互隔离的存储池,通过统一基础设施系统(unified infrastructure system, UIS)实现不同存储池的统一管理,为上层业务提供文件、块和对象3种存储服务,但该文献并未对存储产品的读写性能进行描述,这种融合存储方案可以通过一套存储系统提供多种存储服务,但不同类型的存储服务无法共用硬件设备。文献[12]给出了一种高性能计算(high performance computing, HPC)系统访问采用S3接口的对象存储的实现方案,提出了统一的消息传递接口(message passing interface, MPI)输入输出(input/output, I/O)函数库插件,该插件部署在计算侧(即HPC系统),用于拦截所有的MPI I/O请求,并将其转换为对象访问请求,使HPC系统能够以文件方式访问本地和远端的对象存储。然而,该方式需要在HPC系统的编译阶段添加头文件并链接MPI I/O库函数,增加了实现的复杂性,且要求所有PUT/GET操作均使用字符串流,导致存储访问过程中的数据需要反复进行数据类型转换,影响了存储访问性能。文献[13]提到的技术方案同样需要对计算侧进行改造,开发基于RESTful架构的FirecREST应用程序接口(application program interface, API),以实现HPC系统访问外部对象存储。

目前,国内外通过文件或对象存储结合协议转换网关实现融合存储的产品,存在功能受限、不同类型的存储无法共享硬件设备、性能较差等问题,影响了大模型训练的效率;通过在计算侧部署函数库插件或API实现融合存储的方式,对计算侧具有侵入性,增加了计算侧的复杂性,并影响了存储系



统的访问性能。因此，现有融合存储方案的功能和性能均无法满足大规模智算中心的存储需求。

2.1.2 协议互通融合存储

为了解决协议转换网关带来的存储功能和性能受限，以及计算侧改造带来的对计算侧入侵的问题，本文提出一种融合存储技术，该技术原生支持多种协议，统一元数据管理，不需要进行协议转换，能够避免存储语义损失，且不需要对计算侧进行改造。融合存储技术（协议互通）功能架构如图4所示。

图4中，语义抽象模块将上层不同存储访问协议的访问语义翻译成统一的、数据管理层能识别的访问语义；元数据管理模块负责管理元数据之间、元数据与数据之间的关联关系；数据索引层为所存储的数据建立键值（key-value）索引树，加快数据访问速度；数据持久化层屏蔽底层不同存储设备的差异，将其映射为统一的数据存储单元；增值功能模块负责服务质量（quality of service, QoS）管理、存储配额管理等相关功能；管理模块实现对存储系统的管理。

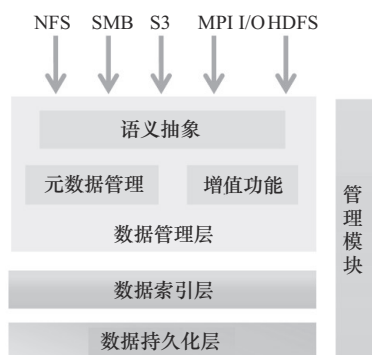


图4 融合存储技术(协议互通)功能架构

协议互通的融合存储技术实现多协议互通的关键是元数据的架构设计，通过一份元数据、一份数据存储支撑多协议无损和无转换互通。统一元数据管理的实现方式如图5所示。

图5中，目录元数据索引同时支持字典序（对象）和数组模式（文件）两种方式访问，有效保障了所有协议的原生语义和性能；innerFile方式支持S3协议的多段上传、多版本特性，且innerFile对其他协议透明，网络附属存储（network attached storage, NAS）协议仅看到多段合并后的

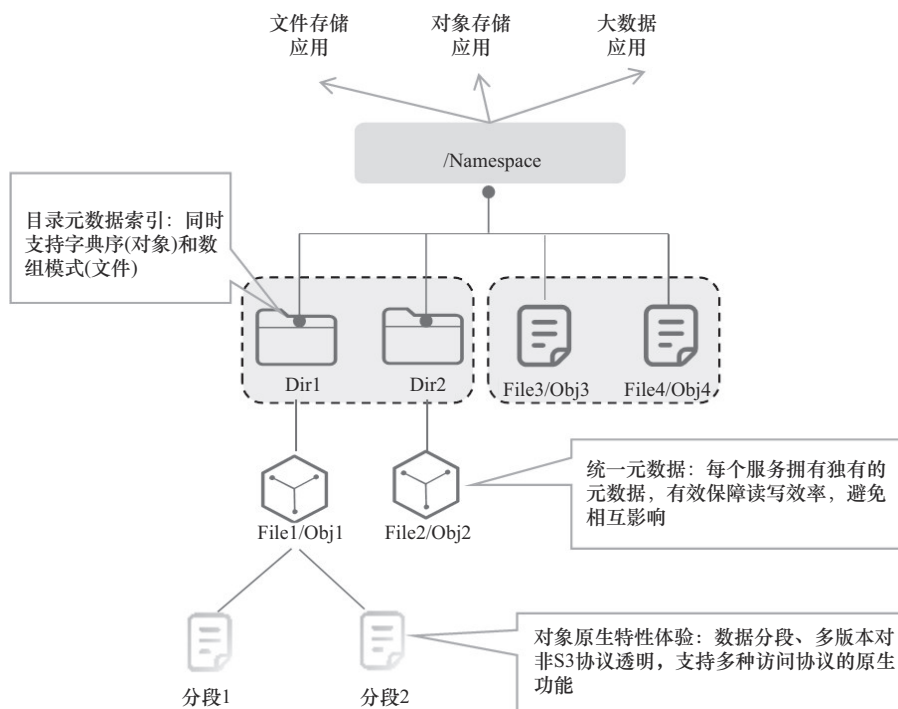


图5 统一元数据管理的实现方式

文件，可以对这些文件进行常规的读、写等操作；同一个文件的公共元数据可以复用，最大限度节约元数据空间占用，不同协议的差异元数据相互独立，避免相互影响，有效保障了读写效率。

协议互通的融合存储技术，支持通用数据和差异化数据读写时的语义互通。采用任意一种协议写入的通用数据可以通过其他协议进行修改和访问，如在申请文件系统时，同步将该文件系统设置桶属性，确保在对象和文件处理通用 I/O 操作时，底层存储系统采用一种方式进行处理。对于差异化数据，本文采用融合的元数据结构，以实现不同协议数据具备相同读写功能的目标，如通过将元数据按文件名进行有序存储，并记录每个目录及其子目录或文件的 position 信息，确保通过 NFS 协议写入的目录和文件在采用 S3 协议读出时能够有序且全路径返回目录、子目录和文件。融合存储实现了多种协议之间的权限互通，支持 NFS/SMB/POSIX 和 S3 跨协议访问权限控制，如 S3 用户和 NFS 用户之间权限互通，需要先配置 S3 用户到 UNIX 用户的映射关系，映射成功后，S3 协议写入的所有文件都会记录正确的用户标识（user identifier, UID）和群组标识（group identifier, GID）等信息，此时其他协议访问时，即可按用户、组、其他（user、group、other, UGO）或 POSIX 访问控制列表（access control list, ACL）进行鉴权。融合存储技术通过统一的分布式锁管理实现不同协议的锁互通，确保任何一种协议一旦加锁成功，其他协议都无法对其进行写操作，从而保证了多协议并发访问时的数据一致性。

千亿模型的训练通常采用万卡规模的智算集群承载，以减少训练时长，训练全流程涉及上百拍字节（petabyte, PB）的文件和对象存储。协议互通的融合存储可以支持训练全流程所需的 S3、NFS、POSIX 等存储协议，在不同阶段实现数据零复制和格式零转换，实现各阶段协同作业的无缝对接和存储空间的高效利用，从而提升训练效率。

2.2 分级存储

智算集群通常同时配置混闪文件存储和全闪文件存储，在满足不同训练阶段读写性能需求的基础上提高投资效益。混闪文件存储用于存储对性能访问要求不高、访问频率不高的数据，全闪文件存储用于存储待训练数据及训练中间状态信息。当前智算集群中的两类文件存储通常独立集群部署，需要用户自己规划数据冷热，并进行数据管理和数据搬迁，消耗更多人力做数据治理。同时，两个存储集群之间的数据复制也需要占用通算和智算服务器的算力。

2.2.1 分级存储当前实现方式

业界现有分级存储技术有服务器内的存储分级和服务器外的存储分级两种。文献[14]提出一种三级存储架构，该架构结合了图形处理器（graphics processing unit, GPU）的高带宽显存（high-bandwidth memory, HBM）、服务器内存、服务器固态硬盘（solid state disk, SSD），旨在解决大规模深度学习广告系统的参数服务器存储问题。该文献主要解决了智算服务器内部的存储分级问题。为了解决基于区块链的大规模工业物联网（industrial Internet of things, IIoT）的存储空间问题，文献[15]提出分级存储架构，通过将区块链的大部分数据存储在云端、少量近期数据存储在区块链中，实现数据的分级存储，但是文中并未提及分级存储是自动分级还是手动分级，且由于数据在两类存储间流动通过普通以太网承载，未对承载网络进行优化，存储分级时的带宽和性能受限。

现有存储分级技术主要面向服务器内部，面向服务器外部的存储分级方式简单、承载网络性能差，不满足智算中心需要的面向智算服务器外部、数据自动、高效分级的需求。

2.2.2 高速自动分级存储

对于智算中心场景，为了提高数据流转效率，降低数据流转对通算和智算服务器的算力消耗，可以在智算集群中引入高速自动分级存储技



术^[16-17]，数据可根据不同的策略在混闪文件存储和全闪文件存储等性能存在差异的存储集群之间实现高速自动流动^[18-19]。数据流动过程不需要用户介入，且不需要消耗通算和智算服务器的算力，前一阶段的输出可以作为后一阶段的输入，实现了训练各阶段任务的无缝对接，达到零等待的效果，显著提升了大模型训练的效率。高速自动分级存储示意图如图6所示。

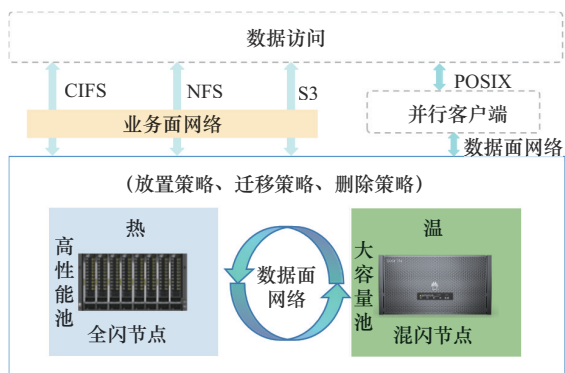


图6 高速自动分级存储示意图

高速自动分级存储主要通过统一命名空间、制定灵活分级策略、搭建高性能网络等满足大模型训练需求。图6中的全闪节点和混闪节点对外提供统一命名空间，用户可以指定数据首次写入时的放置策略，即选择写入混闪文件存储或是直接放在全闪文件存储中，以供智算服务器直接读取；可以提供丰富的数据分级流动策略，如基于文件名、文件类型等文件属性，以及文件存放时长、存储空间使用情况、数据访问频度等，满足策略条件的数据将自动触发流动，迁移过程对业务无感知；可以提供数据预取能力，通过预热策

略来加速计划性任务的冷启动速度，提高训练任务的响应效率；高性能池和大容量池同时接入智算中心的数据面RoCE网络，借助高速无损网络实现冷热数据的高速流动。

高速自动分级存储采用统一命名空间，每个文件只有一个元数据，且位于全闪存储节点上，文件数据则位于全闪存储节点或混闪存储节点中，元数据指向文件数据所在的位置。当数据从全闪向混闪搬迁后，元数据指向同时从全闪改为混闪。由于元数据未变化，搬迁过程中和搬迁完成后都可以使用同一路径访问该数据，上层业务无感，且搬迁可以以较小的数据块为粒度，搬迁过程中的重复数据可忽略不计。数据搬迁时支持文件聚合，小文件写入存储后，以独立的小块数据分开存储，在文件分搬迁过程中，存储系统会将扫描到的同一批小文件聚合成一个较大的块数据进行搬迁，将每秒读写次数（operation per second, OPS）型业务转换为带宽（每秒读写数据量）型业务，提高小文件搬迁效率。分级存储文件聚合示意图如图7所示。

高速自动分级存储在存储数据规模增大时依然可以快速扫描到待搬迁数据。在文件系统创建后，系统会进行一次目录级全量扫描，并生成快照，后续在每次扫描周期到达时，只需要对比数据增量部分的快照与上一次快照之间的差异，即可获得待搬迁的数据，不需要做目录的全量扫描，以确保在大文件规模情况下不会因扫描时间过长而无法完成分级。为了实现数据快速流动、减少数据流动时间，高速自动分级存储将数据流动的

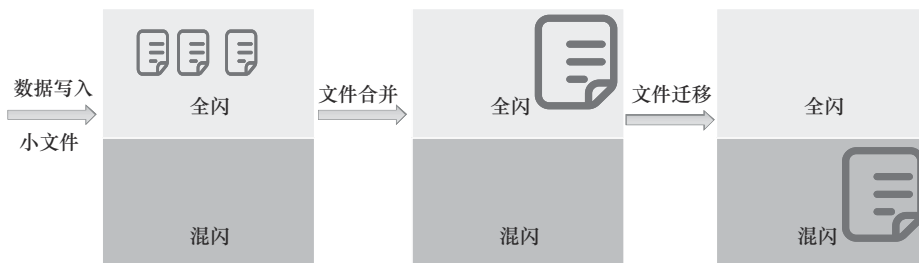


图7 分级存储文件聚合示意图

网络由普通以太网网络改为 RoCE 网络，基于以太网实现 RDMA，利用 RDMA 内核旁路、数据零复制、CPU 卸载的特点实现网络的低时延、高吞吐。

3 存储系统优化验证

3.1 融合存储验证

本文采用 1 750 亿参数规模的大模型对训练过程中的数据存储情况进行验证，对比分析文件对象融合存储和文件对象独立部署时的存储空间占用情况。原始数据类型及规模见表 1。

根据表 1 中的数据类型及现网存储设备使用情况，融合存储与独立部署的空间占用对比分析见表 2。

由于原始数据经过预处理（清洗、标注等）后生成的训练数据集、训练完成后对外发布的模型、用于训练人员分析的部分调试数据等均不需要在文件对象存储系统间复制、重复存储，因此，融合存储可以节省存储空间占用。从表 2 可以看出，采用文件对象融合存储可以节省存储空间 18.02%，且随着原始数据规模增大、训练速度加快，重复存储的数据变多，存储空间节省也随之变大。

3.2 分级存储验证

本文分别采用自然语言处理（natural lan-

guage processing, NLP）和计算机视觉（computer vision, CV）两类模型进行测试验证，对比分析训练数据集准备耗时。分级存储测试组网示意图如图 8 所示。

图 8 中，通算服务器用于训练数据集导入，单台配置 2 个 25 GE 网口，对于图 8 右侧的全闪、混闪独立部署场景，通算服务器同时负责冷热数据迁移。全闪存储配置 6 个存储节点（350 TB），混闪存储配置 3 个存储节点（450 TB），每个存储节点配置 2 个 25 GE 前端网口用于业务数据读写、2 个支持 RoCE 的 25 GE 后端网口用于存储系统后端连接和冷热数据迁移。根据存储系统的性能，数据导入/迁移时，单个全闪存储节点需要 3 台通算服务器达到读写性能上限，单个混闪存储节点需要 2 台通算服务器达到读写性能上限，因此，全闪、混闪分级存储配置 18 台通算服务器，独立部署配置 24 台通算服务器。

在分级存储测试时，全闪和混闪组成一个存储集群，在独立部署测试时分成两个存储集群。首次训练上传数据集时，将分级存储的数据放置策略设置为全闪存储，记录通算服务器上传数据集的时间及训练集预热时间，独立部署侧将数据上传至混闪存储，记录数据集上传时间及迁移至全闪存储（预热）的时间。后续将训练集冷却迁

表 1 原始数据类型及规模

序号	数据类型	大小/TB	备注
1	训练数据集	29	/
2	原始数据	4 982	/
3	checkpoint	3 024	训练 90 天，2 h 一次 checkpoint，单个 checkpoint 大小为 2.8 TB
4	调试数据	336	/
5	模型	0.35	参数量× 2 byte

表 2 融合存储与独立部署的空间占用对比分析

存储方式	对象数据/TB	文件数据/TB	所需总容量/TB	节省百分比
文件对象融合存储	4 982.35（原始数据及模型）	3 389（训练数据集、checkpoint 及调试数据）	8 371.35	18.02%
文件对象独立部署	5 078.55（原始数据、模型、训练数据集及部分调试数据）	5 133.05（训练数据集、checkpoint、调试数据、模型及部分原始数据）	10 211.60	—

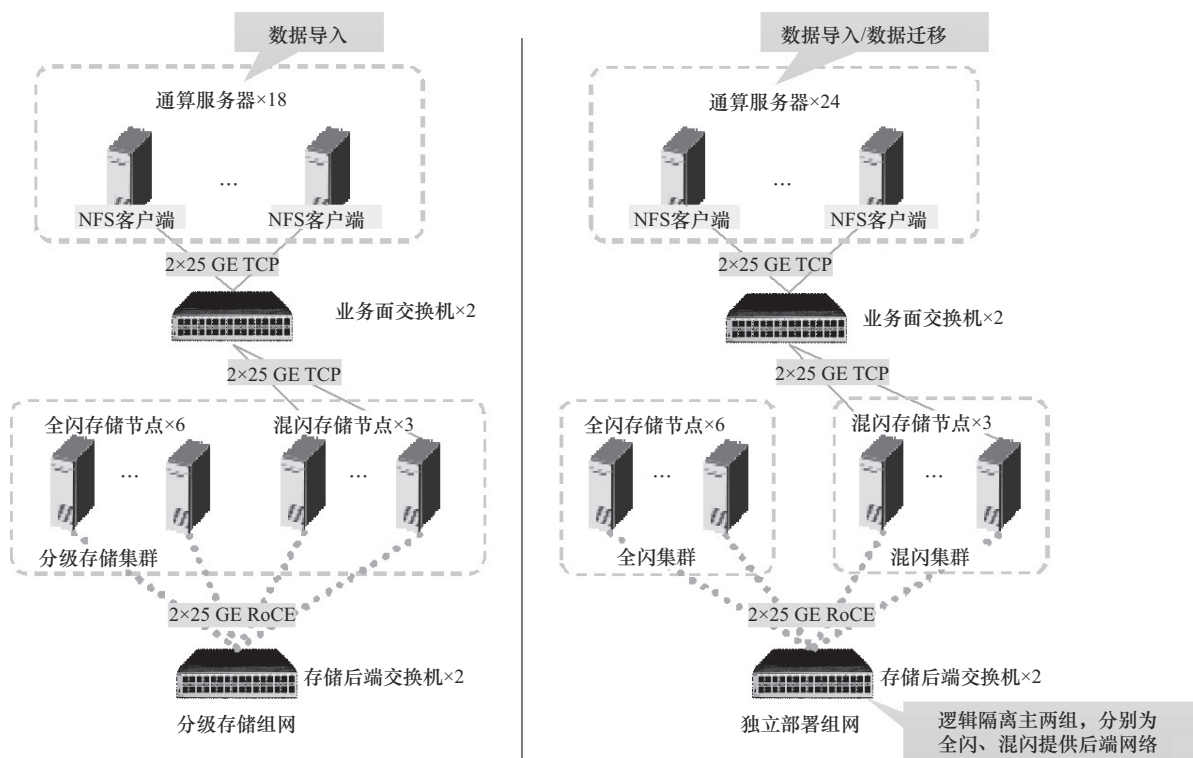


图8 分级存储测试组网示意图

移至混闪存储，测试二次训练时的数据迁移，并记录其时间。

3.2.1 NLP类模型测试结果

NLP类模型测试所用数据集大小为28.1 TB。NLP类模型数据集准备耗时对比见表3。

由表3可知，从用户上传数据集到GPU服务器可以开始训练，全闪混闪自动分级需要耗时1 032 s，全闪混闪独立部署需要耗时19 500 s（11 698 s+7 802 s），分级存储所需耗时约为独立

部署耗时的5.3%，数据集准备时间明显缩短，且不需要专用数据迁移服务器。

3.2.2 CV类模型测试结果

CV类模型测试所用数据集包含1亿个小文件。CV类模型数据集准备耗时对比见表4。

由表4可知，从用户上传数据集到GPU服务器可以开始训练，全闪混闪自动分级需要耗时2 763 s，全闪混闪独立部署需要耗时39 068 s（22 315 s+16 753 s），分级存储所需耗时约为独立部署耗时

表3 NLP类模型数据集准备耗时对比

部署方式	首次训练上传数据集耗时/s	首次训练迁移数据集耗时/s	二次训练迁移数据集耗时/s	数据迁移所需服务器/台
全闪混闪自动分级	1 032	0	7 774	0
全闪混闪独立部署	11 698	7 802	7 785	24

表4 CV类模型数据集准备耗时对比

部署方式	首次训练上传数据集耗时/s	首次训练迁移数据集耗时/s	二次训练迁移数据集耗时/s	数据迁移所需服务器/台
全闪混闪自动分级	2 763	0	2 779	0
全闪混闪独立部署	22 315	16 753	16 655	24

的 7.1%，数据集准备时间明显缩短，且不需要专用数据迁移服务器。

NLP 和 CV 类模型的测试结果中，首次训练上传数据时，自动分级侧的数据通过混闪存储自动流转至全闪存储，数据不会写入混闪存储的硬盘中，由于全闪、混闪存储介质的性能差异，分级存储的数据上传耗时明显缩短。随着数据集规模增加、存储系统规模扩大，自动分级的时间节省效果也更加明显，节省的通算服务器资源也更多。

对不同类型模型测试的结果显示，在大模型训练时，分级存储可以明显缩短数据准备时间，提高数据流转效率，降低数据流转对算力资源的消耗，从而提升大模型训练效率。

4 存储系统优化部署方案

目前，千亿级参数规模的大模型已成为主流，并且逐渐从千亿稠密模型走向万亿稀疏模型^[20]，未来有效样本量和模型参数规模仍会快速增长，大模型训练需要智算集群提供算力更强的计算底座和性能及容量更高的存储底座。融合存储和分级存储可以节省存储容量、减少数据复制、提高训练效率，可以有效支撑当前及未来智算中心的需求。同时，为了降低存储故障的影响范围，可以将存储系统划分为多个存储集群，将存储节点故障范围限制在集群内。智算中心采用的存储系统优化部署方案如图 9 所示。

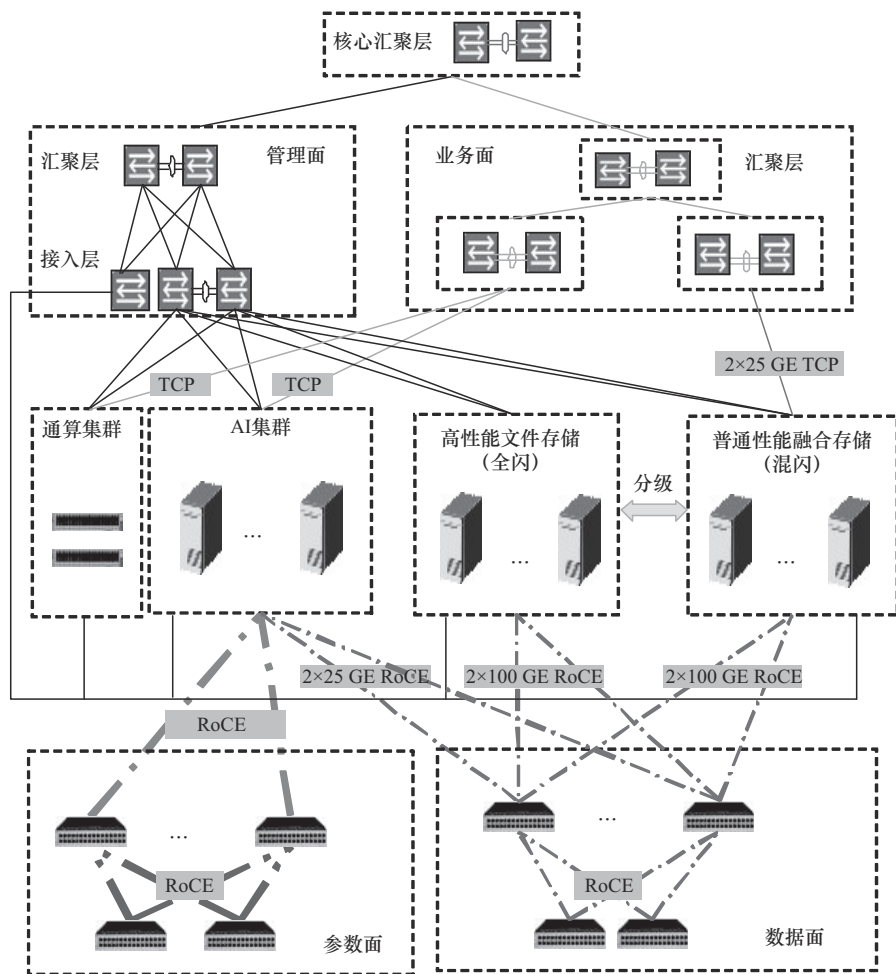


图 9 存储系统优化部署方案



图9中,融合存储通过25 GE端口搭建普通以太网网络(业务/存储面),为用户提供混闪文件存储和对象存储服务,全闪文件存储采用100 GE端口搭建RoCE网络(数据面),为用户提供高性能文件存储服务。同时,融合存储通过100 GE端口接入数据面的RoCE网络,实现全闪文件存储和混闪文件存储的数据自动分级流动。

用户通过25 GE存储面、采用S3协议访问融合存储(对象存储)上传原始数据,通过25 GE存储面、采用NFS协议访问融合存储(混闪文件存储)进行数据预处理,通过100 GE数据面、采用基于POSIX协议的并行客户端访问全闪文件存储进行模型训练。

5 结束语

OpenAI提出的Scaling Law^[21]说明通用大模型的性能依赖于模型规模,规模表现为模型参数量提升和训练数据量提升,这对算力和存储的资源规模提出了更大的需求,万卡及以上规模的智算集群加上百PB及以上规模的融合存储和分级存储成为智算中心建设中需要重点关注的问题。本文提出的融合存储和分级存储可以实现训练全流程数据零复制和格式零转换,可以对外呈现统一的命名空间,提供数据的强一致性,无缝衔接训练的各个阶段,从而显著提升了大模型训练的效率。

参考文献:

- [1] OpenAI. GPT-4 technical report[J]. arXiv preprint, 2023: 2303.08774.
- [2] 曹守欣,赵疏涛,金翊. 基于对象存储的云存储系统研究[J]. 计算机科学与应用, 2014, 4(12): 333-343.
CAO S X, ZHAO L T, JIN Y. The study of cloud storage system based on object-based storage[J]. Computer Science and Application, 2014, 4(12): 333-343.
- [3] 陈曦,朱建涛,何晓斌. 一种面向高性能计算的分布式对象存储系统[J]. 计算机工程, 2017, 43(8): 69-73.
CHEN X, ZHU J T, HE X B. An HPC-oriented distributed object storage system[J]. Computer Engineering, 2017, 43(8): 69-73.
- [4] MA P F, YIN Y S, LAN C, et al. A distributed file system for frequency reading of various file sizes[C]//Proceedings of the 2013 10th Web Information System and Application Conference. Piscataway: IEEE Press, 2013: 339-344.
- [5] 田海东,张明政,常锐,等. 大模型训练技术综述[J]. 中兴通讯技术, 2024, 30(2): 21-28.
TIAN H D, ZHANG M Z, CHANG R, et al. A survey on large model training technologies[J]. ZTE Technology Journal, 2024, 30(2): 21-28.
- [6] 高凯辉,李丹. 数据中心网络性能保障研究综述[J]. 电信科学, 2023, 39(6): 1-21.
GAO K H, LI D. Data center networks with performance guarantee: a survey[J]. Telecommunications Science, 2023, 39(6): 1-21.
- [7] DWIVEDI K, DUBEY S K. Analytical review on hadoop distributed file system[C]//Proceedings of the 2014 5th International Conference-Confluence The Next Generation Information Technology Summit (Confluence). Piscataway: IEEE Press, 2014: 174-181.
- [8] IEEE. IEEE standard for information technology: portable operating system interface (POSIX) base specifications, issue 7: IEEE 1003.1-2008[S]. Piscataway: IEEE Press, 2008.
- [9] LADEKAR S. Converged storage network technologies and guidelines[R]. 2014.
- [10] TANIMURA Y, TAKIZAWA S, OGAWA H, et al. Building and evaluation of cloud storage and datasets services on AI and HPC converged infrastructure[C]//Proceedings of the 2020 IEEE International Conference on Big Data (Big Data). Piscataway: IEEE Press, 2020: 1992-2001.
- [11] H3C. UIS hyper converged unified storage solution[EB]. 2022.
- [12] ARAÚJO DE MEDEIROS D, MARKIDIS S, BO PENG I. LibCOS: enabling converged HPC and cloud data stores with MPI[C]//Proceedings of the International Conference on High Performance Computing in Asia-Pacific Region. New York: ACM Press, 2023: 106-116.
- [13] CRUZ F A, DABIN A J, DORSCH J P, et al. FirecREST: a RESTful API to HPC systems[C]//Proceedings of the 2020 IEEE/ACM International Workshop on Interoperability of Supercomputing and Cloud Technologies (SuperCompCloud). Piscataway: IEEE Press, 2020: 21-26.
- [14] ZHAO W J, XIE D P, JIA R L, et al. Distributed hierarchical GPU parameter server for massive scale deep learning ads systems[J]. arXiv preprint, 2020: 2003.05622.
- [15] WANG G, SHI Z J, NIXON M, et al. ChainSplitter: towards blockchain-based industrial IoT architecture for supporting hierarchical storage[C]//Proceedings of the 2019 IEEE Interna-

tional Conference on Blockchain (Blockchain). Piscataway: IEEE Press, 2019: 166-175.

- [16] 曹纪磊. 基于网络存储系统中虚拟存储分级存储技术的研究[J]. 软件, 2022, 43(11): 83-87.

CAO J L. Research on hierarchical storage technology based on virtual storage in network storage system[J]. Software, 2022, 43(11): 83-87.

- [17] 赵晓南, 李战怀, 曾雷杰, 等. 分级存储管理技术研究[J]. 计算机研究与发展, 2011, 48(S1): 105-111.

ZHAO X N, LI Z H, ZENG L J, et al. Research on hierarchical storage management technology[J]. Journal of Computer Research and Development, 2011, 48(S1): 105-111.

- [18] ZHANG T R, HELLANDER A, TOOR S. Efficient hierarchical storage management empowered by reinforcement learning extended abstract[C]//Proceedings of the 2023 IEEE 39th International Conference on Data Engineering (ICDE). Piscataway: IEEE Press, 2023: 3869-3870.

- [19] BU K, WANG M, NIE H S, et al. The optimization of the hierarchical storage system based on the hybrid SSD technology[C]//Proceedings of the 2012 Second International Conference on Intelligent System Design and Engineering Application. Piscataway: IEEE Press, 2012: 1323-1326.

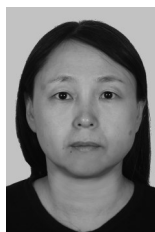
- [20] FEDUS W, ZOPH B, SHAZEER N. Switch transformers: scaling to trillion parameter models with simple and efficient sparsity[J]. arXiv preprint, 2021: 2101.03961.

- [21] KAPLAN J, MCCANDLISH S, HENIGHAN T, et al. Scaling laws for neural language models[J]. arXiv preprint, 2020: 2001.08361.

[作者简介]



曹原铭 (1983-), 男, 中国移动通信集团设计院有限公司高级工程师, 主要研究方向为云计算、人工智能算力基础设施等。



雷鸣 (1973-), 女, 中国移动通信集团设计院有限公司高级工程师, 主要研究方向为云计算、人工智能算力基础设施等。



刘芹 (1977-), 女, 中国移动通信集团设计院有限公司正高级工程师, 主要从事与云计算、人工智能等算力资源相关的咨询及研究工作。



牛瑛霞 (1971-), 女, 现就职于中国移动通信集团设计院有限公司, 主要研究方向为算力网络、人工智能、数据业务平台等。



武振宇 (1975-), 男, 中国移动通信集团设计院有限公司高级工程师, 主要研究方向为云计算、算力网络、智算中心等。



潘洁 (1978-), 女, 中国移动通信集团设计院有限公司高级工程师, 主要研究方向为算力网络安全和网络信息安全。