

引用本文:张轲,刘梦涵,高梓文,等.基于算力网络的分布式机器学习通信优化方法[J].光通信技术,2026,50(3):41-45.

基于算力网络的分布式机器学习通信优化方法

张轲¹,刘梦涵¹,高梓文¹,周勇¹,吴莹²,彭振文^{2*}

(1.湖南大学信息科学与工程学院,长沙 410082;2.雁城区块链研究院,湖南 衡阳 421000)

摘要:针对基于算力网络的分布式机器学习中计算节点间长通信距离导致通信成本高、效率低的问题,提出一种分布式机器学习通信优化方法。该方法在梯度压缩阶段结合随机梯度量化与Elias Gamma编码减少单次通信量,在梯度合并阶段采用合并梯度无等待反向传播(MG-WFBP)策略降低通信频率。通过模拟2~1 024个节点的算力网络环境,对GoogleNet、ResNet-50和mnistNet模型进行训练测试。结果表明:当节点数扩展至1 024个时,该方法的加速效果明显优于传统MG-WFBP方法;在合理压缩率下,该方法能显著缩短仅用于通信的时间,从而提升训练效率。

关键词:算力网络;机器学习;梯度量化;分布式系统

中图分类号:TN256 文献标志码:A 文章编号:1002-5561(2026)03-0041-05

DOI:10.13921/j.cnki.issn1002-5561.2026.03.007

Communication optimization method for distributed machine learning based on computing power networks

ZHANG Ke¹, LIU Menghan¹, GAO Ziwen¹, ZHOU Yong¹, WU Ying², PENG Zhenwen^{2*}

(1. College of Computer Science and Electronics Engineering, Hunan University, Changsha 410082, China;

2. Yancheng Blockchain Research Institute, Hengyang Hunan 421000, China)

Abstract: To address the high communication cost and low efficiency caused by long inter-node distances in distributed machine learning over computing power networks, This paper proposes a communication optimization method for distributed machine learning. In the gradient compression phase, the method combines stochastic gradient quantization with Elias Gamma coding to reduce the per-communication data volume. in the gradient merging phase, it adopts the merged gradient wait-free backpropagation MG-WFBP strategy to decrease the communication frequency. Simulations are carried out on a computing power network environment with 2 to 1 024 nodes, training and testing GoogleNet, ResNet-50, and mnistNet models. The results show that when the number of nodes scales to 1 024, the proposed method achieves significantly better speedup than the traditional MG-WFBP method. Moreover, under a reasonable compression ratio, the method can markedly reduce the time spent solely on communication, thereby improving training efficiency.

Key words: computing power networks, machine learning, gradient quantization, distributed system

0 引言

随着机器学习技术的不断创新和算力需求的不断增加,基于大规模数据挖掘的机器学习已成为智能

发展的重要手段。分布式机器学习通过多个计算节点协同工作,基于网络通信进行参数更新与模型训练,从而加速机器学习过程^[1-4],可以有效解决数据集中存储带来的隐私问题。但是,考虑到计算量大、训练数据量大、模型规模大这三方面因素,分布式机器学习的实施面临的主要挑战是提高计算节点之间的通信效率^[5]。算力网络^[6-8]作为一种资源分配和管理机制,可以更高效地管理数据流,减少冗余传输,从而降低延迟,提高带宽利用率,便于大量数据在不同节点之间进行传输以实现模型的共同训练^[9-10]。基于算力网络的分布式机器学习能够作为一种保护数据隐私、

收稿日期:2024-05-20。

基金项目:长沙市科技计划项目(kh2301011)资助。

作者简介:张轲(1986—),男,四川宜宾人,博士研究生,高级工程师,从事算力网络和人工智能应用科研工作,先后主持和参与国家及省部级课题10余项,软件著作权及发明专利20余项。

*通信作者:彭振文(1984—),男,湖南祁东人,湖南大学博士研究生,研究方向为算力网络计量。



张轲,刘梦涵,高梓文,等:基于算力网络的分布式机器学习通信优化方法

协调大规模参与方的学习手段,可以在不共享终端原始私密数据的情况下进行模型学习^[11],但是在该学习过程中,长距离通信存在通信成本较高的问题。

为了解决此问题,Zhang Hao等人^[12]基于深度神经网络(DNN)的分层结构,提出在计算底层更新的同时反向传播顶层的参数,实现了通信和计算重叠的算法^[13-14],称为无等待反向传播(WFBP)。Yang You等人^[15]提出了一种单层通信机制(SyncEASGD),它将所有梯度合并到每次迭代结束时的单个全局归约操作中,此方法可以减少大部分数据通信的启动时间。考虑到将一些短的通信任务合并为单个任务可以减少整体通信时间,Shi Shaohuai等人^[16]提出了一种合并梯度WFBP(MG-WFBP)的分布式训练算法,尽可能避免传输少量数据,进一步优化了分布式机器学习通信效率。但以上几种通信优化方式应用到基于算力网络的分布式机器学习时都存在网络带宽低,从而无法获得较好性能的问题。基于此,本文提出一种基于算力网络的分布式机器学习通信优化方法。

1 通信优化方法设计

本文提出了一种基于算力网络的分布式机器学习通信优化方法,其核心架构包含2个主要模块:梯度压缩模块和梯度合并模块。在训练开始前根据梯度压缩模块预设的压缩率确定网络模型每层通信量,进而确认梯度合并方案。然后使用MG-WFBP方法将多个梯度进行合并,并将合并后的梯度传输到中央服务器。以上设计通过采用梯度压缩模块和梯度合并模块分别减少了通信量,提升了通信频率,从而降低了基于算力网络的分布式机器学习所需的通信开销,提高了整体的效率和性能。

1.1 梯度压缩模块

梯度压缩模块包含两部分:第一部分采用随机梯度量化方法将梯度进行阈值为 s 的量化操作;第二部分使用Elias Gamma编码方法将转化后的梯度进行高效编码。

1.1.1 随机梯度量化方法

随机梯度量化方法将机器学习模型中需要通信的梯度进行压缩,以减小梯度的存储空间和传输量。这一过程将浮点数表示的梯度转换为较低比特率的整数或定点数表示形式,从而减少存储需求。

随机梯度量化方法中的量化方程如式(1)所示。

$$Q_s(v_i) = \|v\|_2 \operatorname{sgn}(v_i) \varepsilon_i(v, s) \quad (1)$$

式中: v 为量化梯度,阈值 s 为大于等于1的超参数(对应量化级别的个数), $\operatorname{sgn}(v_i)$ 是符号函数,其表达式如式(2)所示。

$$\operatorname{sgn}(x) = \begin{cases} -1, & \text{if } x < 0 \\ 0, & \text{if } x = 0 \\ +1, & \text{if } x > 0 \end{cases} \quad (2)$$

$\varepsilon_i(v, s)$ 是独立随机变量,其表达式为

$$\varepsilon_i(v, s) = \begin{cases} l/s, & \text{概率为 } 1 - p(|v_i|/\|v\|_2, s) \\ (l+1)/s, & \text{其他情况} \end{cases} \quad (3)$$

式中:概率函数 $p(\alpha, s) = \alpha s - l, \alpha \in [0, 1]; l$ 为层数。

1.1.2 Elias Gamma编码方法

Elias Gamma编码使用可变长度编码来表示整数序列,用于将非负整数序列编码为二进制字符串。该编码的基本思想是使用一个前缀码,其中每个编码对应一个正整数,编码的长度根据数字的大小而变化,较大的数字使用更长的编码,这种编码方法能对整数序列进行高效的压缩^[17]。

编码的完整流程如算法1所示。对于一个整数 n ,首先根据采用floor函数计算出前缀码长度 L ,然后将该整数转换为二进制作为后缀码 b 并与前缀拼接,得到Elias Gamma编码^[18]。

算法1 Elias Gamma编码算法

输入:待编码的正整数 n

输出:Elias Gamma编码后的二进制字符串

- 1) $b \leftarrow n$ 转换为二进制形式
- 2) $L \leftarrow \operatorname{floor}(\log_2 n) + 1$ //计算二进制编码的位数
- 3) if $n == 1$ then
- 4) return b
- 5) else
- 6) 前缀 $\leftarrow$ 使用 $L-1$ 个0
- 7) result $\leftarrow$ 前缀+ L 的二进制+ b
- 8) return result
- 9) end if

在解码时,需要按照相反的步骤进行操作:对于编码 N ,从左到右读取编码,首先找到 N 前缀中连续的0的个数,然后根据这个数量确定整数的位数。接下来,读取剩余位数作为整数的二进制表示,并将其还原解析为十进制形式 n 。

1.2 梯度合并模块

在梯度合并模块中,本文使用MG-WFBP方法,该

方法引入了合并通信的思想,能够有效降低算力网络环境下分布式模型训练所需的通信频率,从而大幅度提高分布式机器学习的训练效率。

在训练开始前,该方法首先从模型的最后一层向前进行遍历。对于每一对相邻的 $l-1$ 层和 l 层,该方法比较了将2层通信合并进行和分开进行这两者之间的通信代价,可表示为

$$\tau_b[l-1] + t_b[l-1] < \tau_c[l] + a \quad (4)$$

式中: $\tau_b[l-1]$ 表示第 l 层开始计算梯度的时间戳, $t_b[l-1]$ 表示每次迭代中反向传播的时间, $\tau_c[l]$ 表示第 l 层开始梯度通信的时间戳, a 表示全归约(all-reduce)的启动时间(延迟)。若式(4)成立,则说明将第 $l-1$ 层与 l 层的通信进行合并的收益更大,能够节约通信成本,所以执行合并操作。

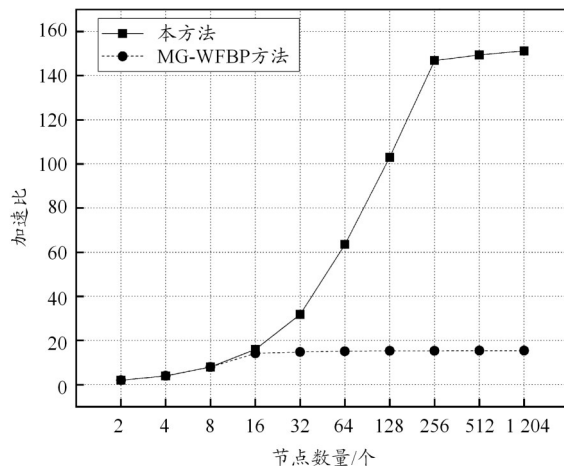
2 实验设置和结果分析

为验证通信优化方法的有效性,本文进行了对比试验。由于硬件限制,算力网络场景的模拟根据实际的单图形处理单元(GPU)性能和网络性能模型进行。

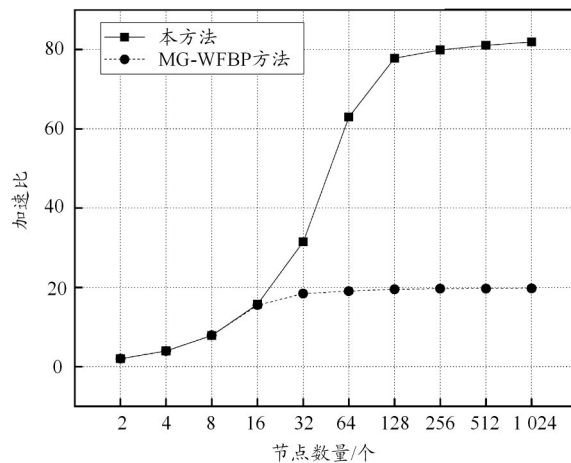
本文实验从2个节点扩展到1 024个节点来训练GoogleNet、ResNet-50和mnistNet模型,分别模拟了这3种模型在算力网络场景中使用MG-WFBP方法和本文提出的通信优化方法(下文简称本方法)时的性能指标(以加速比表征),结果如图1所示。其中,mnistNet不是特定的机器学习模型名称,而是在处理基于MNIST数据集的手写数字识别任务时编写的神经网络。这3种模型兼顾了计算量和通信量较大的情况,能够较全面地反应本文提出的通信优化方法的正面效果。

从图1(a)和图1(b)可以看出,在2~16个节点的集群中,本方法与MG-WFBP方法的性能表现基本一致。随着节点数量从16个扩展到1 024个,本方法带来的加速效果明显优于MG-WFBP方法。从图1(c)可以看出,在节点数量扩展到128个时,本方的加速效果才优于MG-WFBP方法。这是因为MNIST数据集中每个图片的数量较小,并且mnistNet模型较为简单,其通信量比GoogleNet和ResNet-50模型少。

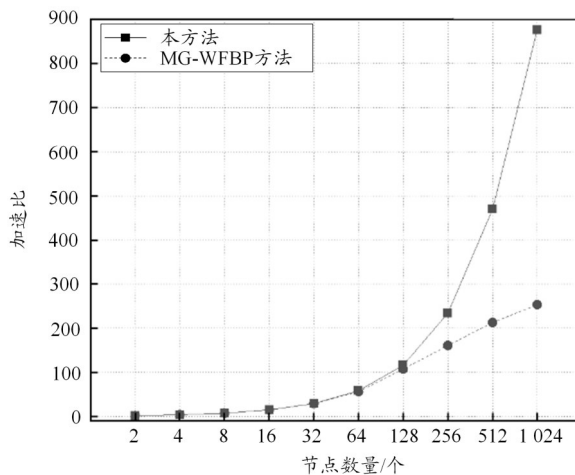
另外,在节点数量较少的情况下,本方法在性能提升方面并不显著。这是因为当节点数量较少时,每个节点需要处理的任务相对较多,通信开销在整个训练过程中的占比相对较小。因此,即使通过通信优化减少了通信时间,但对总的训练时间的减少有限,导



(a) ResNet-50模型



(b) GoogleNet模型



(c) mnistNet模型

图1 3种模型在算力网络场景中采用不同方法的性能指标对比图
致性能提升不明显。并且分布式训练的性能提升很大程度上依赖于并行度。节点数量较少意味着并行度有限,这限制了本方法能够获得的加速比。随着节

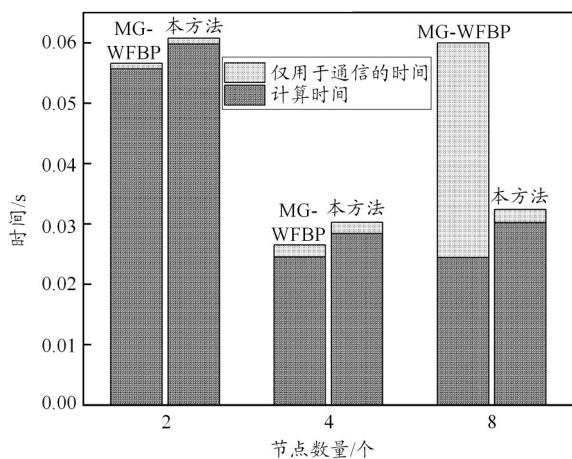
张轲,刘梦涵,高梓文,等:基于算力网络的分布式机器学习通信优化方法

点数量的增加,并行度提高,通信优化的效果才更加显著。

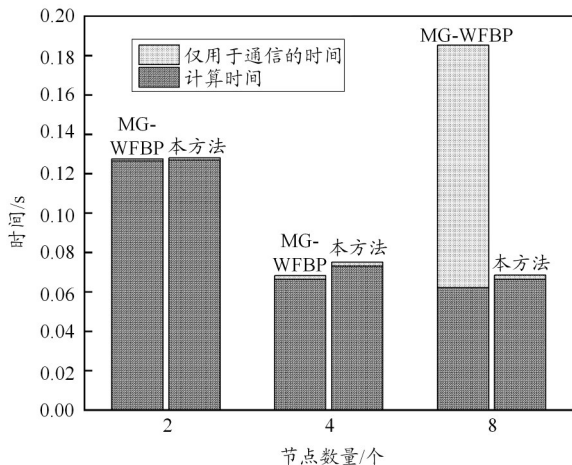
为了对比本文方法与 MG-WFBP 方法在通信量上的优化效果,本文采用通信时间反映通信量,在基于算力网络的分布式机器学习过程中比较了 GoogleNet、ResNet-50 和 mnistNet 模型分别使用这 2 种方法进行前向传播与反向传播的计算时间与仅用于通信的时间,结果如图 2 所示。

从图 2(a)和图 2(b)可以看出,针对 GoogleNet 和 ResNet-50 这 2 种模型,本方法仅用于通信的时间低于 MG-WFBP 方法。其中,节点数为 2 和 4 时均为小规模情况,因此通信时间优化量少,而节点数为 8 时引入了大量的通信与同步开销,本方法通过对梯度参数进行量化压缩,因此有效地减少了仅用于通信的时间,从而降低了每次迭代的总时长,进一步提升了在算力网络场景下的分布式机器学习训练效率。

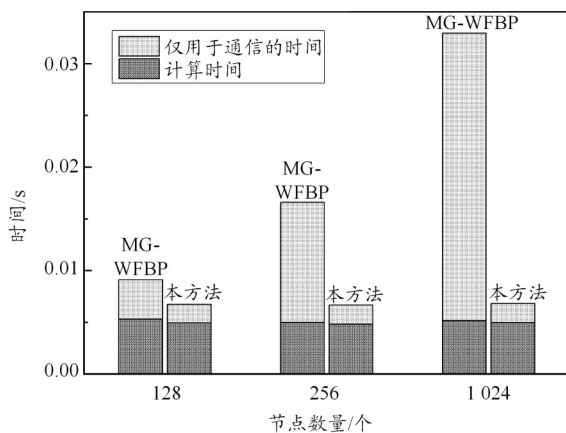
从图 2(c)可以看出,在计算时间相同的情况下,



(a) ResNet-50 模型



(b) GoogleNet 模型



(c) mnistNet 模型

图 2 3 种模型使用不同方法时的计算时间与仅用于通信的时间对比图

本方法节省了大量仅用于通信的时间。这是因为 mnistNet 模型尺寸较小,本文选择了较大集群规模,以测试本方法性能。

以上实验结果证明了本方法在通信量与训练效率方面的优越性。说明本方法可以更好地利用计算资源,提高分布式机器学习的效率和可扩展性。这对于解决大规模分布式系统和边缘设备面临的通信带宽限制和计算资源有限的问题具有重要意义。

3 结束语

本文提出了一种基于算力网络的分布式机器学习通信优化方法,通过梯度压缩模块和梯度合并模块来减少通信量和通信频率。在梯度压缩模块中使用随机梯度量化和 Elias Gamma 编码方式,有效减少了通信量。在梯度合并模块中引入了 MG-WFBP 方法,可以有效降低通信频率。通过对 2~1 024 个节点的算力网络场景模拟实验结果表明,本文提出的通信压缩模块与编码方式能够有效压缩多节点间的通信量,从而大幅减少仅用于通信的时间,提升训练效率。因此,该方法可以在基于算力网络的分布式机器学习场景中应用,以加快训练过程中各节点之间的数据交换进程,提高整体效率和性能。

参考文献:

- [1] Xing E, Ho Q R, Xie Pengtao, et al. Strategies and principles of distributed machine learning on big data[J]. Engineering, 2016, 2(2): 179-195.
- [2] Lee K, Lam M, Pedarsani R, et al. Speeding up distributed machine learning using codes[J]. IEEE Transactions on Information Theory, 2017, 64(3): 1514-1529.
- [3] Duchi J, Hazan E, Singer Y. Adaptive subgradient methods for online

learning and stochastic optimization[J]. Journal of Machine Learning Research, 2011, 12(7): 2121–2159.

[4] 肖雄,唐卓,肖斌,等.联邦学习的隐私保护与安全防御研究综述[J].计算机学报,2023,46(5):1019–1044.

[5] Xing E, Ho Q R, Dai Wei, et al. Petuum: a new platform for distributed machine learning on big data[J]. IEEE Transactions on Big Data, 2015, 1(2): 49–67.

[6] 黄光平,罗鉴,周建峰.算力网络架构与场景分析[J].信息通信技术,2020,14(4):16–22.

[7] 贾庆民,丁瑞,刘辉,等.算力网络研究进展综述[J].网络与信息安全学报,2021,7(5):1–12.

[8] 李少鹤,李泰新,周旭.算力网络:以网络为中心的融合资源供给[J].中兴通讯技术,2021,27(3):29–34.

[9] Xiao Xiong, Li Chuanying, Jiang Bingting, et al. Adaptive search strategy based chemical reaction optimization scheme for task scheduling in discrete multiphysical coupling applications[J]. Applied Soft Computing, 2022, 121: 108748.

[10] Yang Li, Xiao Xiong, Zhang Xuedong, et al. A real-time partition generation mechanism for data skew mitigation in spark computing environment[J]. Journal of Grid Computing, 2023, 21(4): 62.

[11] Xiao Xiong, Tang Zhuo, Yang Li, et al. FDSFL: Filtering defense strategies towards targeted poisoning attacks in IIoT-based federated learning networking system[J]. IEEE Network Magazine, 2023, 37(4): 153–160.

[12] Zhang Hao, Zheng Zeyu, Xu Shizhen, et al. Poseidon: an efficient

张轲,刘梦涵,高梓文,等:基于算力网络的分布式机器学习通信优化方法

communication architecture for distributed deep learning on GPU clusters [C]// ATC. Proceedings of the 2017 USENIX Conference on USENIX Annual Technical Conference ATC. Santa Clara: ATC, 2017: 181–193.

[13] Awan A A, Hamidouche K, Hashmi J M, et al. S-caffe: Co-designing mpi runtimes and caffe for scalable deep learning on modern gpu clusters [J]. Acm Sigplan Notices, 2017, 52(8): 193–205.

[14] Shi Shaohuai, Wang Qiang, Chu Xiaowen, et al. A dag model of synchronous stochastic gradient descent in distributed deep learning [C]// IEEE. Proceedings of 2018 IEEE 24th International Conference on Parallel and Distributed Systems (ICPADS). Singapore: IEEE, 2018: 425–432.

[15] Yang You, Buluç A, Demmel J. Scaling deep learning on GPU and knights landing clusters [C]// IEEE. Proceedings of SC17: International Conference for High Performance Computing, Networking, Storage and Analysis. USA: IEEE, 2017: 1–12.

[16] Shi Shaohuai, Chu Xiaowen, Li Bo. MG-WFBP: Efficient data communication for distributed synchronous SGD algorithms [J]. IEEE Transactions on Parallel and Distributed Systems, 2021, 32(8): 1903–1917.

[17] 白福均.云环境下全文索引压缩关键技术研究[D].贵阳:贵州大学,2019.

[18] Fenwick P M. Punctured Elias codes for variable-length coding of the integers [R]. Auckland: Department of Computer Science, The University of Auckland, 1996.