

TensorCIM: Digital Computing-in-Memory Tensor Processor With Multichip-Module-Based Architecture for Beyond-NN Acceleration

Yiqi Wang¹, Graduate Student Member, IEEE, Zihan Wu, Weiwei Wu², Leibo Liu², Senior Member, IEEE, Yang Hu², Member, IEEE, Shaojun Wei², Fellow, IEEE, Fengbin Tu², Member, IEEE, and Shouyi Yin², Senior Member, IEEE

for beyond-NN
(GCN, DLRM),

spG gather algebra consumption

Redundancy-Eliminated Gathering Manager 消除冗余, 访问对齐

MCM-CIM

Tensor CIM

Abstract—While neural networks (NNs) have achieved great results in various intelligent tasks like image classification and speech recognition, real-world scenarios have more applications beyond just NN processing like graph convolutional network (GCN) and deep-learning recommendation model (DLRM), which typically consist of sparse gathering (SpG) and sparse algebra (SpA). Their large application size leads to substantial data movement. Although the fusion of digital computing-in-memory (CIM) and multichip-module (MCM) can reduce data movement efficiently and scale out CIM's capacity in a high-yield solution, the MCM-CIM system raises new challenges for beyond-NN acceleration: SpG involves repeated off-chip DRAM access, interchiplet access, and redundant reduction operations; SpA suffers from inter-CIM workload imbalance and intra-CIM under-utilization. Thus, we design TensorCIM as the CIM processor chiplet with three corresponding features: 1) the redundancy-eliminated gathering manager (REGM) dynamically maintains frequently accessed features and reduction results in the CIM to eliminate redundant accesses and reductions; 2) the equal operation-based CIM initializer (EOCI) calculates effective multiply-accumulation (MAC) operations and initializes CIM macros with a balanced inter-CIM workload at the subarray level; and 3) the input-lookahead CIM (ILA-CIM) architecture looks ahead at future inputs to fully utilize CIM logic. The fabricated MCM-CIM system consumes only 3.7 nJ/Gather for the GCN model, achieving 8.3-TFLOPS/W algebra efficiency at FP32.

arch-optim: MCM-CIM for data trans. & workload.

Index Terms—Computing-in-memory (CIM), floating-point, graph neural network (NN), multichip-module (MCM), recommendation system, sparsity.

energy effective: dynamic cache, workload, pre-implementation, REGM, EOCI, ILA-CIM

Manuscript received 17 January 2024; revised 8 April 2024; accepted 24 May 2024. Date of publication 13 June 2024; date of current version 30 January 2025. This article was approved by Associate Editor Priyanka Raina. This work was supported in part by the National Key Research and Development Program under Grant 2018YFB2202600, in part by the NSFC under Grant U19B2041 and Grant 61774094, in part by the Beijing Science and Technology (S&T) Project under Grant Z191100007519016, and in part by the Beijing Innovation Center for Future Chip. (Corresponding authors: Fengbin Tu; Shouyi Yin.)

Yiqi Wang, Zihan Wu, Weiwei Wu, Leibo Liu, Yang Hu, Shaojun Wei, and Shouyi Yin are with the School of Integrated Circuits, Beijing Innovation Center for Future Chip, and Beijing National Research Center for Information Science and Technology, Tsinghua University, Beijing 100084, China (e-mail: yinsy@tsinghua.edu.cn).

Fengbin Tu is with the Department of Electronic and Computer Engineering, The Hong Kong University of Science and Technology, Hong Kong, China (e-mail: fengbintu@ust.hk).

Color versions of one or more figures in this article are available at <https://doi.org/10.1109/JSSC.2024.3406569>.

Digital Object Identifier 10.1109/JSSC.2024.3406569

Equal Operation-based CIM Initializer 均衡的运算初始化

WHILE neural networks (NNs) have achieved great results in various intelligent tasks [10], real-world scenarios have more applications that have computational and data-movement requirements beyond those seen in typical NN processing, like graph convolutional network (GCN) [8] for graph-based learning and deep-learning recommendation model (DLRM) [14] for recommendation (see Fig. 1). GCN aggregates features on a large graph and then performs an NN-based combination. DLRM gathers and reduces items from an embedding table and generates recommendation results via fully connected layers. However, prior artificial intelligence (AI) chip research focuses on NN-only applications, and we demand a new architecture for beyond-NN acceleration.

To realize efficient beyond-NN acceleration, we first analyze the workload features of beyond-NN applications. Beyond-NN applications typically consist of sparse gathering (SpG) and sparse algebra (SpA). SpG gathers and reduces tensors from sparsely distributed addresses, like GCN's aggregation phase and DLRM's embedding layer. SpA refers to NN-based sparse tensor multiplication for the gathered tensors (GTs) like GCN's combination phase and DLRM's fully-connected (FC) layer. Besides, due to the large application size, data movement is the main bottleneck for beyond-NN acceleration, and FP32 precision is usually required for high accuracy.

Our solution for beyond-NN applications is the fusion of digital computing-in-memory (CIM) and multichip-module (MCM), which reduces data movement while maintaining high precision and accommodates large application sizes, respectively. In the past few years, digital CIM has proven to be an efficient and precise architecture for reducing data movement [4], [5], [13], [16], [19], [20]. The large-scale beyond-NN applications bring huge computational and storage requirements and motivate the demand for scaling out digital CIM processors. MCM provides a high-yield scalable solution for CIM scaling by integrating multiple chiplets in one package to build an MCM-CIM system [15], [28], [29].

Fig. 2 shows a typical MCM-CIM system with four CIM chiplets. In our MCM-CIM system, workloads are partitioned to each chiplet. The CIM chiplet mainly relies on its own DRAM channel to perform SpG and SpA. The MCM architecture enlarges CIM capacity. Intermediate data are

GCN: Graphic Convolutional Network

通过聚合邻居节点的特征, 更新每个节点, 的表示, 提取局部信息

flow: ① Aggregation ② combination $H^{(l+1)} = \sigma(\hat{A} \cdot H^{(l)} W^{(l)})$
聚合信息 与自身信息结合 + 非线性变换

特点: 稀疏, 局部性, 图结构 (2. 规则与无规则), 精度要求

DLRM Deep Learning Recommendation Model

用于个性化推荐, 融合稀疏特征 (用户点击历史, 商品 ID) ...
稠密特征 (年龄, 性别) ...

flow: ① Embedding: 嵌入层, 稀疏特征映射为低维向量

② Interaction: 特征交互, Embedding vector $\otimes$ 稠密特征 $\rightarrow$ 组合特征

③ MLP: 全连接网络, 输入为组合特征, 输出为行为

特点: Embedding 层巨大, gather + reduce 频繁, 稀疏计算 挑战

sparsity Gather: GCN 邻接聚合, DLRM: 嵌入向量查找

sparsity Algebra: 稀疏向量的乘积运算

spG: 频繁访问 DRAM 和 chiplet, 带宽压力大
相同特征的重复处理

spA: 计算负载不均

CIM 空闲 $\rightarrow$ 利用率低

Tensor CIM: Redundancy-Eliminated-Gathering Manager

spG: 建立 dynamic cache

动态维护索引特征与中间结果于 CIM 内

避免重复读取/计算

Equal-Operation CIM Initializer

spA: 基于 GT 和 WT 的非零次及乘计算有效 MAC

并行激活不同 CIM 核子阵列

CIM 有效 MAC 差过大

Input-Look-Ahead CIM

spA 计算时, LLA 预读输入, 并行激活子阵列

每次输入只有部分阵列激活

sparsity FP32 accuracy spG spA access-unordered

data movement

$\rightarrow$ digital CIM + MCM

MAC inside of SRAM + multi chiplet package as system

CIM accuracy

independent
connection, DRAM channel, storage + expander

3 challenge and resolution

A: 对于不规则的数据, 需先 spG + spA MAC + accuracy
gather + reduce.

B: CIM: effective, high accn. expansion.

在 SRAM 单元中加入 2 输入 NOR 门 $\rightarrow$ 按位乘.

阵列周边加数子加法器, partial sum.

C. MCM: 一个芯片存储能力不够, 需要多个组合.

D. spG spA 问题

Tensor CIM architecture

主要结构:

① ILA - CIM Core ($\times 8$).

负责主计算单元

② Redundancy Eliminated Gathering Manager ($\times 8$)

③ Equal Operation - CIM Initializer

CIM 初始化器

④ Input buffer ⑤ Global buffer

⑥ SIMD Core ⑦ Top Controller.

3 种数据流: spG: CIM gathering
spA init: CIM Initializer
spA Compute: CIM Computation.

flow: ① spG model.

load tensor from DRAM.

REGM: select frequently-ass. storage into CIM

CIM Core implement feature gathering $\rightarrow$ Gathered Tensor.

GT $\rightarrow$ global buffer.

② spA Init.

EOCI load weight Tensor from DRAM, analysis sparsity arch.

distribute WT to CIM core's SRAM

read GT from Input Buffer.

③ spA Compute:

ILA-CIM 预读, 激活子阵列

FP32: MAC.

result store $\rightarrow$ Global buffer

SIMD: 激活.

datapath 复用, spG spA 优化.

Redundancy Eliminated Gathering Manager mechanism.

spG: Beyond-NV 第一步: 根据不规则的访问流量 Access Tensor
从 feature tensor 中提取相关特征.
进行归约, 生成: Gathered Tensor.

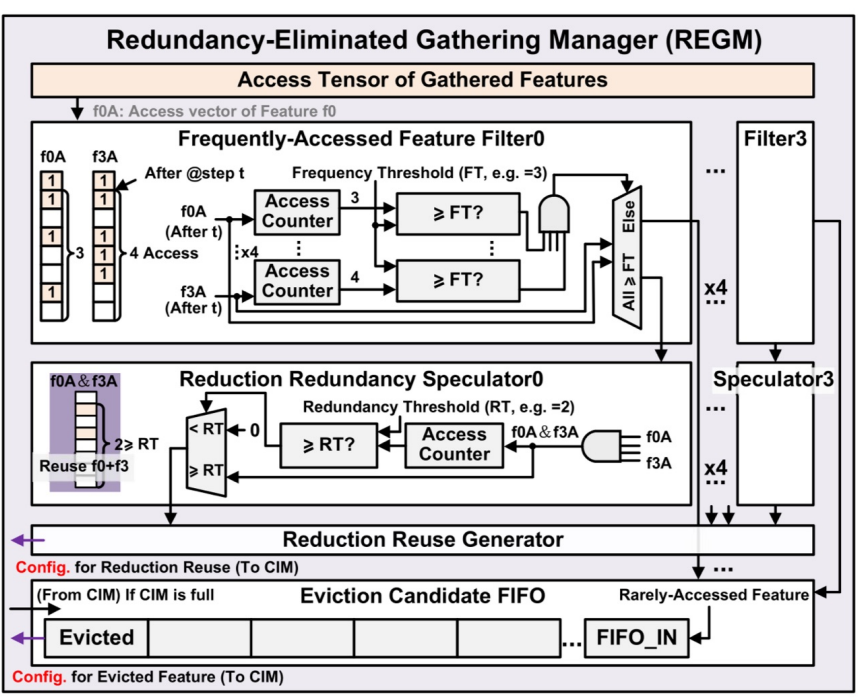
在 GCN 中: spG: AT: 图的邻接矩阵.

DLRM: AT: 用户历史交互
数据情况

访问模式高度集中且重复: 大量节点访问同一节点/局部性
↳ 缓存命中率高, 带宽是否浪费

- 优化策略: ① Access Tensor reorder: 让相似访问模式的操作集中处理.
② 将 spG 工作分配给不同的 chiplet. 每个 chiplet 访问局部性 DRAM.
③ REGM 冗余去除:
记录访问频率高的特征及其归约结果
CIM 中缓存以重用
避免重复访问 重复计算, 减少 DRAM 带宽和 chiplet 间的通信.

architecture.



- ①/② 获取 AT 信息.
统计每个特征的归约访问次数.
与设置的阈值作比较.
↓
CIM $\in CF$
③ 对于 $f_0 + f_3$ 使用冗余推测器.
获取未来可能的访问次数.
↓
CIM
生成 configuration 指导 CIM 写入

①. Frequently-Accessed Feature Filter (x4)
判断哪些是高频访问 AT, 并保存在 CIM 中.

②. Reduction Redundancy Speculator (x4)
判断哪些归约结果值得存入缓存
生成存储与重用控制信号 指导 CIM 存取行为

③. Reduction Reuse Generator ④. Eviction Candidate FIFO

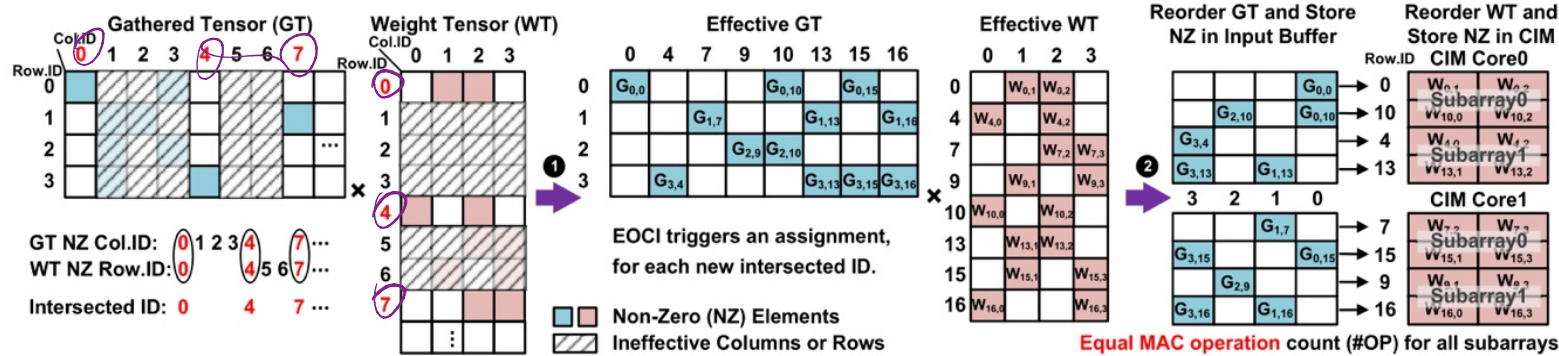


Fig. 12. Equal operation-based CIM initialization example: ① effective OP recognition and ② balanced WT/GT assignment.

Equal Operation CIM Initializer:

spA中不同的CIM工作负载不同 (稀疏性)

在 CIM 初始化, 通过智能分配 Weight Tensor 和 Gathered Tensor 实现负载均衡

How: ① 获取非0索引, 求交集 (有效计算只在两者交集)

② 对于有效的ID, 统计对应行列的计算数...

③ 选择合适的 CIM 子阵列映射, 平衡工作负载

④ 生成 CIM 的权重矩阵, WT Tiling.

→ 找出有效计算 → 形成有效计算的WT GT → 软件分配计算 → 划分数据 → CIM 计算

architecture:

Effective Operation Recognizer

检测 GT 和 WT 的非零 ID 交集.

On-Demand Weight Fetcher

从 DRAM 中加载真正被使用的 WT.

Balanced Inter-CIM Distributor

管理各个 CIM 子阵列负载/工作分配

How:

从 global buffer 读取 GT

↓ E. ID

GT FIFO

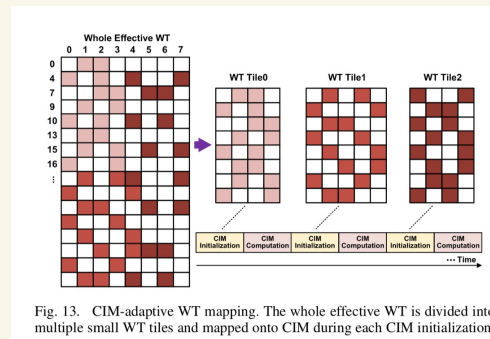


Fig. 13. CIM-adaptive WT mapping. The whole effective WT is divided into multiple small WT tiles and mapped onto CIM during each CIM initialization.

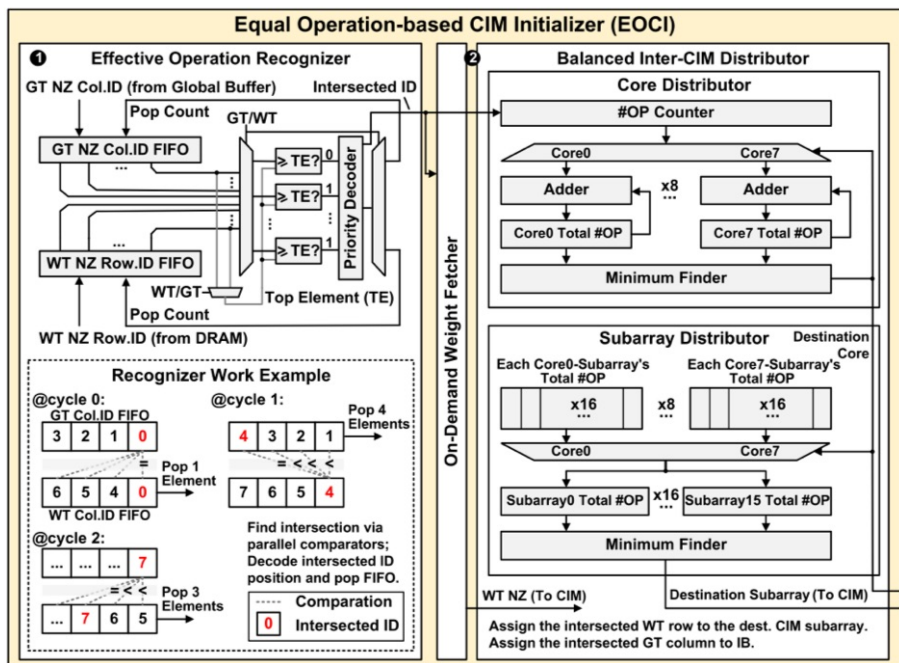


Fig. 14. EOCI architecture: ① effective OP recognition and ② balanced WT/GT assignment.

从 DRAM 中读取 WT

↓ EID

WT FIFO

↓ ↓
并行比较器

收集

WT FIFO - 所有 GT

GT FIFO - 所有 WT

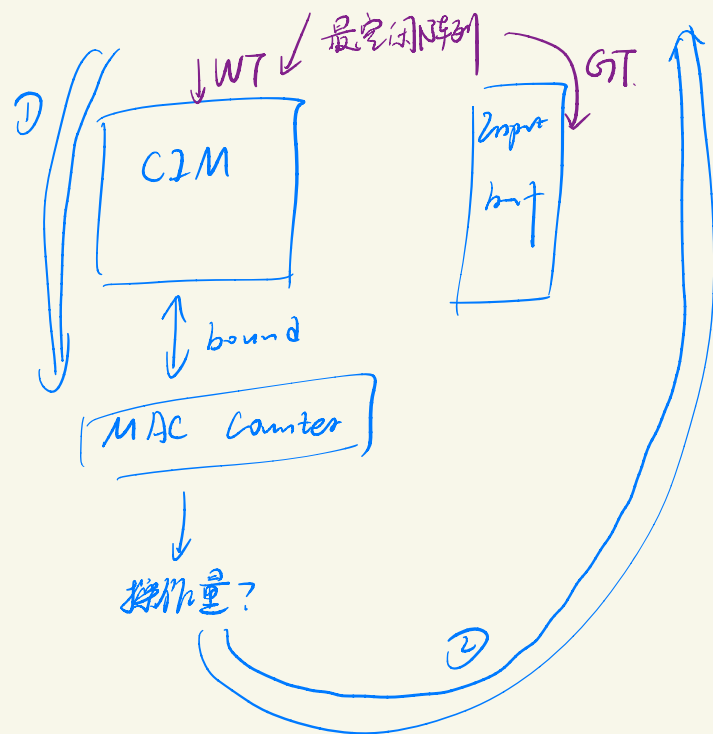


每以别一个 Effective ID.

Input Buffer ← load from Global Buffer GT (col)

On-Demand - Fetcher ← load from DRAM WT (row)

Balanced Inter-CM Distributor.



Input - Look Ahead CIM

负载均衡 EOC1 后, 仍可能存在 空闲计算单元存在.

传统方式 每次只处理一行输入, 只能激活部分阵列

→ MAC 闲置 / 延时

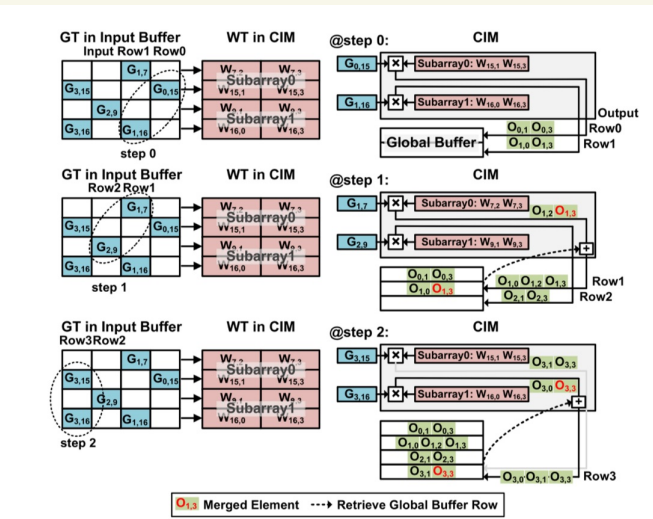


Fig. 16. ILA-CIM computation example.

传统方式 每次只处理一行: for Input Row 0 → 只会激活 subarray 0.

此处优化: step 0: 处理 G0,15 时同时预看 G1,10, 把 subarray 1 也激活, 输出 Row 0 和 Row 1 写入 global buffer.

step 1: global buffer 中已有部分和 先取出 O1,0 与结果作稀疏累加

steps: --

architecture: 流水线结构.

①. Input Look Ahead: 提前准备行输入, 填满每个子阵列 FIFO

②. SRAM CIM Macro: { SRAM Bank, Subarray Multiplier 计算 (Booth signals).

③. Output Merges: 合并子阵列结果并写入 global buffer

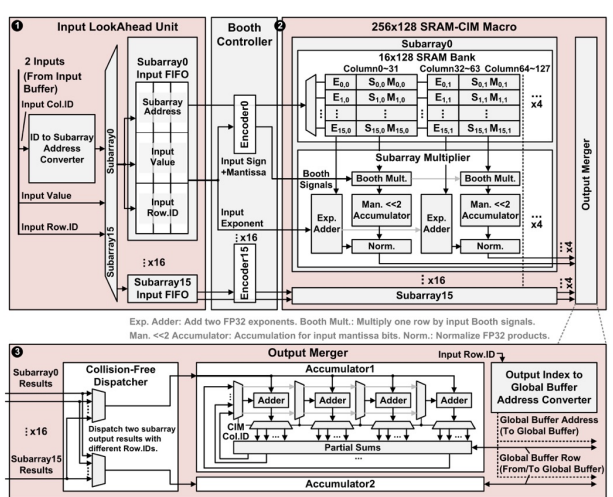


Fig. 18. ILA-CIM architecture: ① input lookahead, ② CIM multiplication, and ③ output merging.

How: ① 每个 CLM Core 有 4 个 2LA-CLM macro
每个 2LA-CLM macro 有 16 个 subarray
所有 2LA-CLM macro 共享 Input Look Ahead unit / Booth Controller.
每周读取 2 个输入 ID + Value.
并推入 FIFO.
当准备就绪后可同时激活 16 个子阵列

② CLM 乘法:

每个 subarray 有: 16×128 个 6T SRAM, 存储权重 + 4 个 Booth 的选率法器.

③ 输出合并:

所有子阵列的输出结果传向 Output Merger.

映射到 global buffer 地址.

把中间结果 $\hookrightarrow$ 读回并累加.

最终结果写回 global buffer.

2LA-CLM: ^{的 datapath} 用于 spA, 也可用于 spG Reduction.