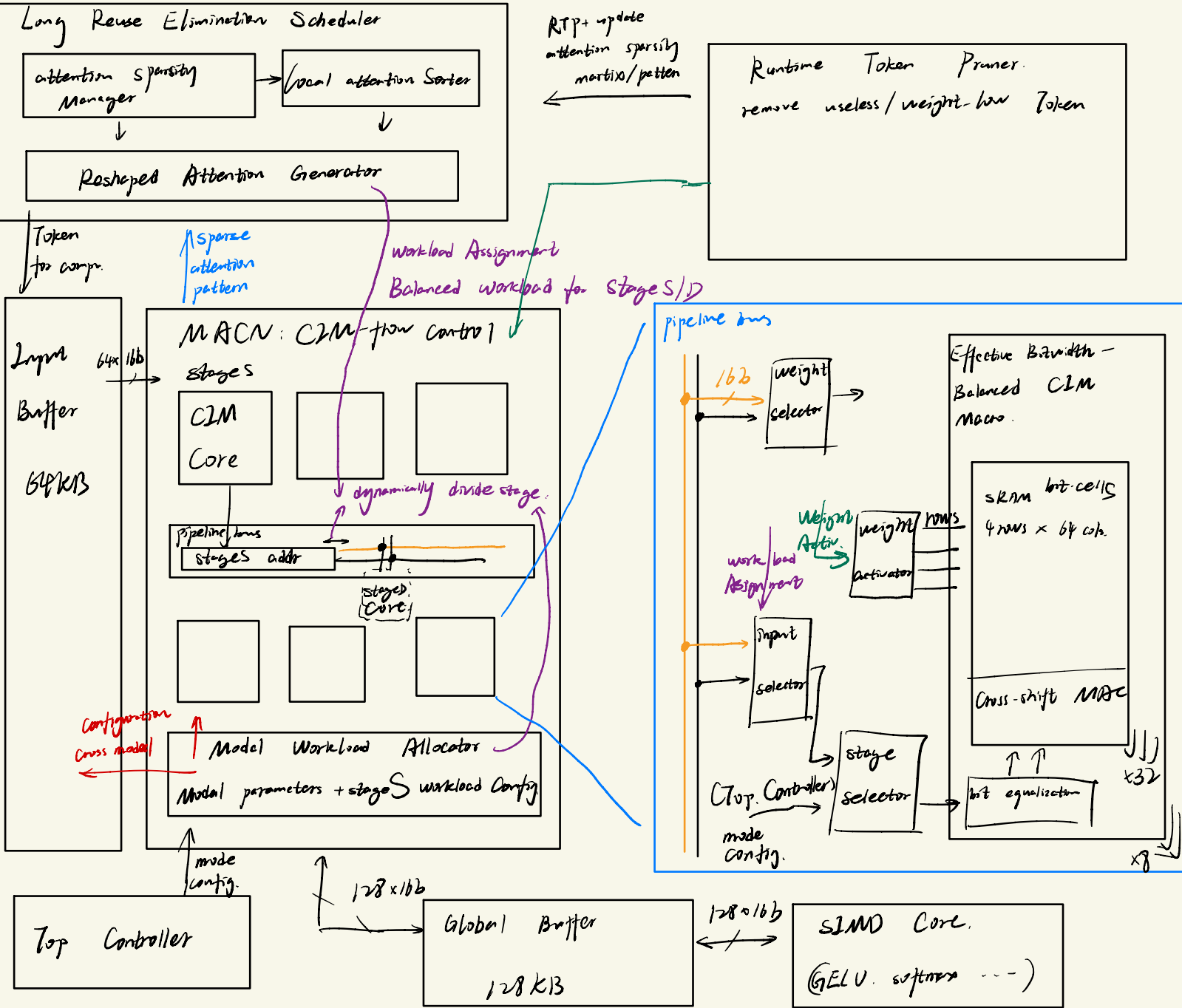


CIM-Core: stages / stageD: Dynamichly Distro, generate Q.KV / Q.KT. A.V



Long Reuse Elimination Scheduler: 管理 attention sparsity
控制 attention 的 token 输入顺序和 sparsity matrix 稠度

RTP 数据 info → LRES 数据 → MACN
Input Buffer stageS (Q.KV) stageD (attention 计算) 前 接收 Input Buffer 和 MACN 稠度方式
stageD Q.KT. A.V SIMD 处理

Runtime Token Pruner: 实时将 Token 的 weight/activate 数据 减少无用 token

EBB-CIM: 调节不同的 位稀疏性, 提高 CIM 单元计算效率
动态调节 位置, 适应不同模型 Token ...

Cross-model attention

Model-Workload-Allocator: 切换 Model. (FC-Attention)

SRAM bit cells, 存储 weight.
cross shift MAC: 动态控制 数据流向通道.
bit equalization: 平衡不同层位的 bit 分布.

计算结果存入 Global buffer, SIMD 完成后续计算能力.

LRES: stage D: $A \cdot V$, $A = Q \cdot K^T$, $Q = \vec{q}$ (input) $K = \vec{k}$ (weight in CIM)
 一些大长期闲置不运算, (k 只向 q 算, q 只向 k 算) $\rightarrow$ 与 CIM 资源何不用 short
 LRES: original attention $\xrightarrow{\text{reshape}}$ global-like + local-like $\left\{ \begin{array}{l} \text{global like: 提取对多数数据相关的 } q, k \text{ 保留更长时间, } \rightarrow \text{long reuse?} \\ \text{local like: } q, k \text{ 相对少数 (相邻) 数据相关, } \rightarrow \text{quick refresh 减少占用} \end{array} \right.$

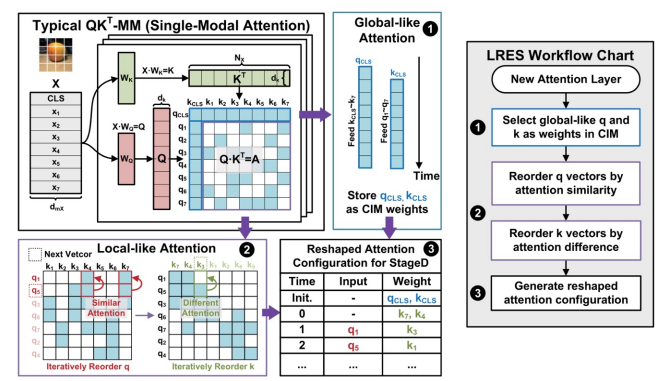


Fig. 6. Long reuse elimination workflow, with a QK^T -MM example.

How: ① 从 attention matrix 中找出最常用, 与其他 Token 交互通量的量
 放在 CIM 中长期使用 (cache hit)
 ② 重排列表 q, k , 尽可能组成相似块
 使用相似的 q 在一起, 与之相关联 k 在一起
 ③ generate reshaped attention configuration
 $\rightarrow q$ 之间可被输入 / CIM 中并行载入 k

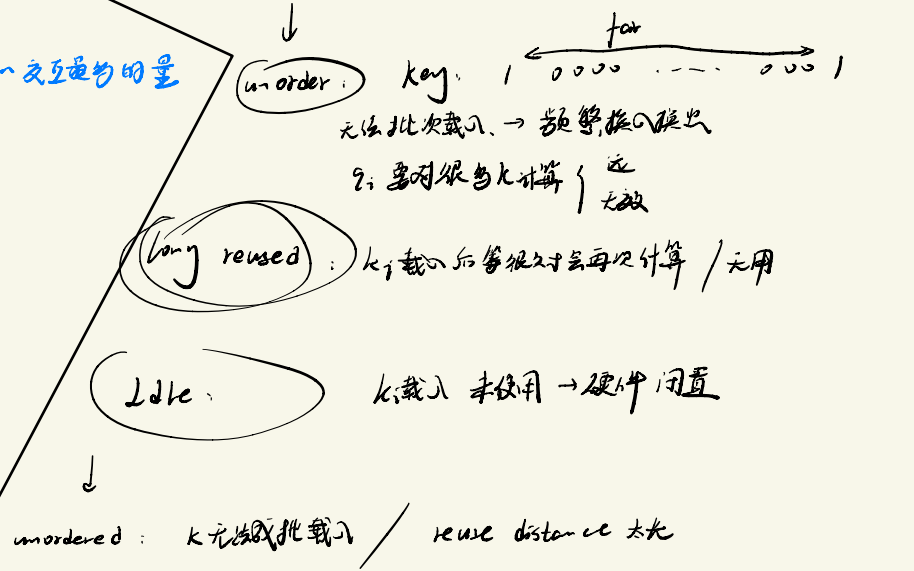
long idle: long reused: global-like.

attention sparsity: local-like.

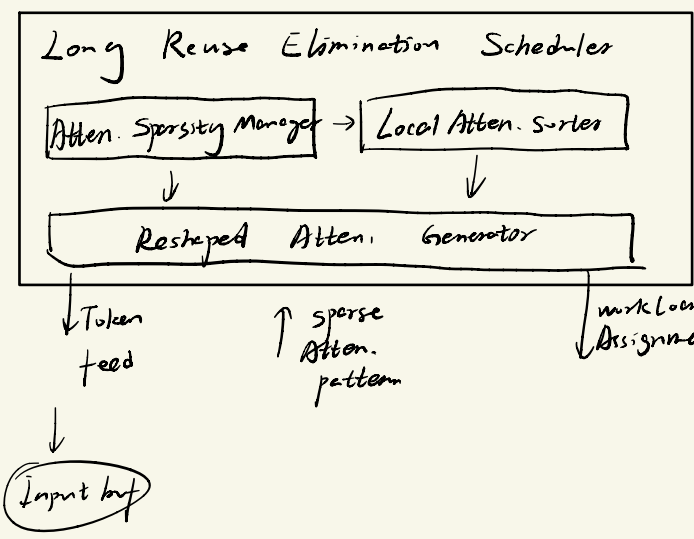
comput./transport chaos: Configuration.

Original unordered Attention matrix?
 in transformer (google): $A = Q \cdot K^T \rightarrow$ Attention matrix.
 $A[i, j] \rightarrow$ i th Token (Query i) 对 j th Token (key j)
 的关注程度.

原则上: Token 可以 dense, 但实际上稀疏明显



LRES architecture.



step ①: Attention sparsity Manager.
 Select: global like q, k
 step ②: Local Attention Sorter
 re-order q, k (removed) $\xrightarrow{\text{format}}$ local-like.
 step ③: Reshaped Attention Generator
 根据 global-local 排成 输入/runtime configuration

ASM: $\rightarrow$ 稀疏的 + RTP (remove-info).

①: row sum: $\rightarrow$ global q (找出与最多 Token 相关的 query) 发出 attention
 column sum: $\rightarrow$ global k 哪个 key 被更多的 query 关注

思考题① 经过网络 → global → FIFO

↓:



LAS: 尽可能以 attention 相似的设计去, 让它们共享一组 k (core-window)

并行等可能, slide-window 聚合.

Attention Similarity-based Clustering

使多个 query 共享少量 key CIM 硬件友好, 数据搬运损耗 ↓

RAG: 生成实际的 attention 注意力配置表.

time: ? Input: ? weight: ? global/local ?

attention similarity-based clustering

相似 attention pattern 的 token 归为一类.

简化 / 共享 attention 计算 → effective + accelerate + compression.

→ 很多 token 关注的东西是一样的 / 相似.

事先分析 query 向量特征 + cluster + 每个 cluster 保留一个代表 query.

compression: { memory: 少存
bandwidth: 无需对所有计算

cluster → global + local

剪掉掉不重要的 attention block.

global-like / less query? / reduce useless computation.

Reshaped Attention Generator

↓ Configuration

1. 控制 q, k 生成顺序 stages.

2. 指定 CIM cores 读取什么 q, k staged.

slide-window 聚合. (CNR?)

让 query 只关注附近的 token.

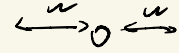
减少计算量和 memory reuse distance.

初始的 Attention 是 dense 的

$$A = Q \cdot K^T \quad (512 \times 512)$$

60% query 只关注 window=1 的, 即前后最多关注 1 个.

每次只计算 $2w+1$



CNR fixed weight. / Attention slide windows
runtime compute. weights.

fixed pattern / pattern based on segment

2D piz. / language token.

? sparsity optimisation.

slide-window: { reduce data transportation.
keep k used frequently
parallel + stream pace.

Local Attention Sorter

(mechanism)

rows: 让下一个输入时 q 与上一个有最大的 similarity.

两个 q attention 的 k 相似.
无需替换 CIM 中的 weights.

col: 调整 k 的排列型.
出去不用回来

FIFO: ① global - q, k .

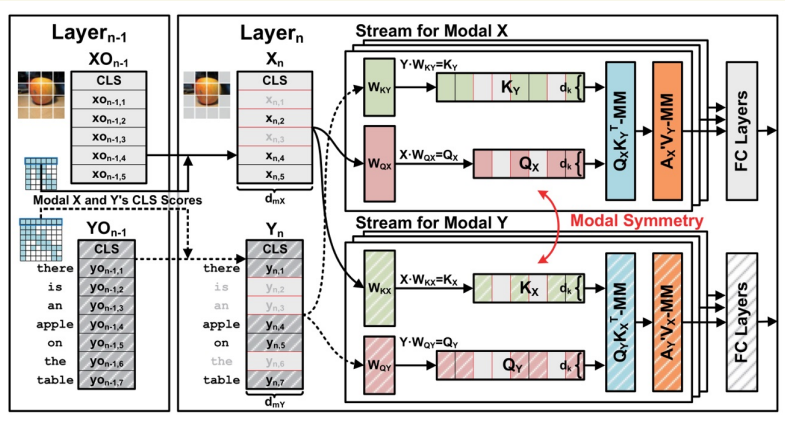
② local q

③ local k .

RTP: ① Token prune based on Class Score (CLS)

Transformer中有专门处理整体语义的CLS Token

CLS Token 对Token的 attention高, 对Token重要.
 并根据CLS进行裁剪.



flow: layer₁中有CLS token.
 attention matrix 第一行为CLS Token对其他Token的 attention.
 得分低的是被裁剪的Token
 (裁剪后可能导致 model X 和 model Y Token 数量不等).

challenge: 2: (MACN解决)
 X的Token比Y少, 裁剪后更少.
 在 cross-modal attention中, Q_x 小于 K_y 或
 Q_x 会空转等待 K_y 计算
 ↓
 pipeline stages 延时, CIM利用率↓

MACN arch.
 Multiply-Accumulate CIM Network.
 包括: ① Modal Workload Allocator

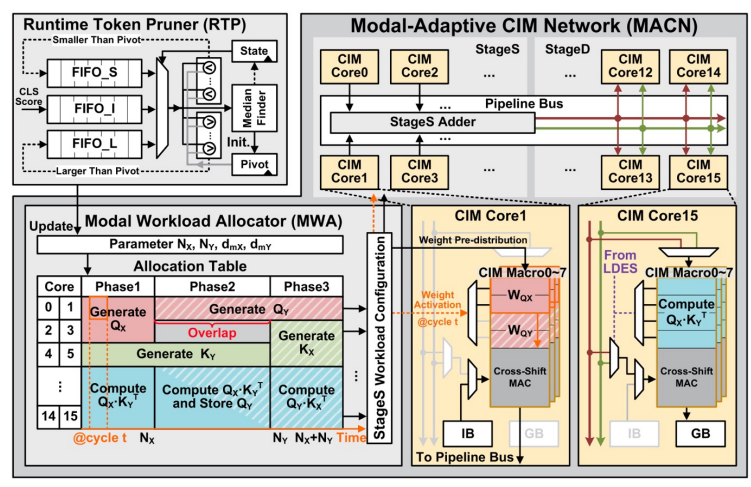


Fig. 11. RTP and MACN architecture.

CLS Token是如何生成的?
 Classification.
 是一个学习的 输出向量

Runtime Token Pruner
 mechanism.
 从上层得到CLS
 利用 median-of-medians top-n 演算法.
 计算最重要的n个 token.

median-of-medians top-n.
 选出无冗余数据中第n/前n大的数.
 如要找前K大的数.

 过程: pivot

对所有数, 分为3部分:
 Larger than pivot Equal than pivot Smaller than pivot
 如: $k \leq |L|$, 就等于在 'L' 中找前k大的数
 $|L|+1 < k \leq |L|+|E|$, 结果为 pivot
 $k > |L|+|E|$, 在 'S' 中找前 $k - |L| - |E|$ 大的数.

②. 16个 CIM cores. ③. pipeline bus.
 MACN有个模式: FC: CIM core 并行运行 (parallel)
 (pipeline). attention: MWA根据RTP分配 stage S/2

challenge 2 解法: (X, Y token 不同长度, 负载不均导致 Q_x 空转).

重新计算两模块的 Q/k , 对称及错!

How: phase 1: Model X stream 开始,

core 0-3: Gene. Q_x

core 4-5: Gene. K_T

core 6-15: compute. $Q_x \cdot K_T^T$

Q_x 先完成计算

phase 2: Symmetric Model Overlap.

Q_x 完成, core 0-5: 生成 Gene. Q_T

在 Gene. K_T 的同时 生成 Q_T

phase 3:

core 4-5: 生成 K_x

core 6-15: compute $Q_T \cdot K_x^T$

$Q_x \cdot K_x / Q_T \cdot K_T$ 生成时间大致相同

bit serial 计算 (位文为列乘法)
 在 INT8 模式下, 每个数据为 8 位, 要 8 个 cycle 完成计算

EB Balanced:
 How: 如有: I_0, I_1, I_2, I_3
 $EB-I_0 = 4, EB-I_1 = 1, \dots$
 解法: 将 I_3 的高位搬到 I_1 上, 一起计算

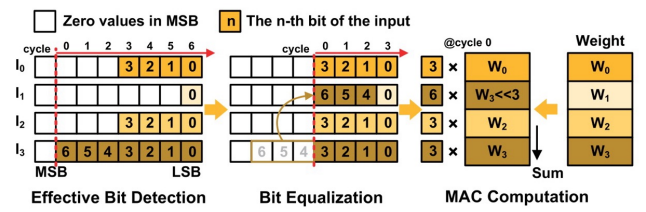


Fig. 13. EB-balanced multiplication.

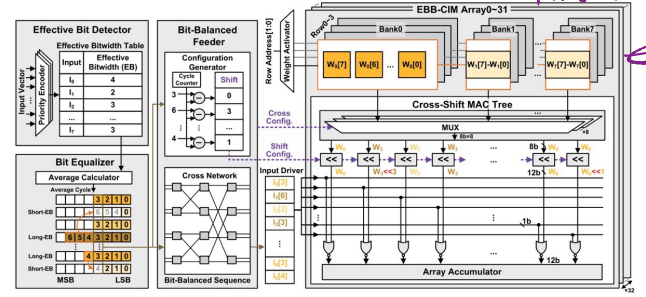


Fig. 14. EBB-CIM macro architecture.

如 0 位计算结果: $I_0[3] \cdot W_0 + I_3[6] \cdot (W_3 \ll 3) + I_2[3] \cdot W_2 + I_3[3] \cdot W_3$
 让每个 cycle 最大化利用可用位。

- EB Balanced mechanism:
- Components:
 - 32 个 EBB-CIM arrays, 每个数组输入信号
 - Effective Bit Detector: 检测有效位数
 - Bit Equalizer: 比对有效位数并左移进行 balance
 - Bit-Balanced-Feeder: 发送 (MSB - LSB) 控制 MAC tree.

EBB array:
 每个 array 有: ① 4 rows, 64 cols, 分 8 个 bank.
 ② cross-shift MAC tree 进行权重相加和移位运算。
 ③ 4 个 8 列的 CIM 逻辑。

Effective Bit Detector:
 输入 8 个元素 (一个 vector)
 用 priority Encoder 判断最高位 $\rightarrow$ EB.

Bit Equalizer:
 长位元分配给短位元
 Bit Balance Feeder
 从 MSB 到 LSB 的发送。

signal
 { Cross-configuration: 权重值在 MAC tree 中交换
 Shift-configuration: 权重左移位数。
 同时
 MAC: 读取 8 个 weight 数据 cross/shift configuration 相加
 每个 bank 每行 8 个