



北京大学
PEKING UNIVERSITY

人工智能的硬件基石

从物理器件到计算架构

第十一讲：AI芯片设计-III

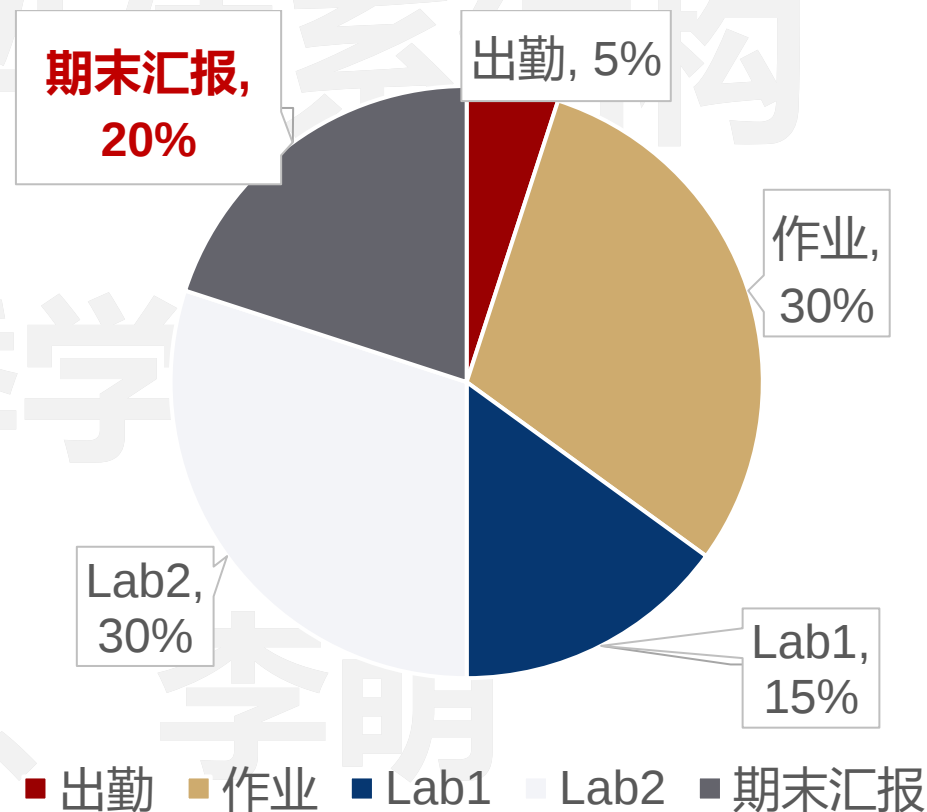
主讲：陶耀宇、李萌

2025年春季

注意事项

• 课程作业情况

- 第5次作业时间：5月13号-5月27号（2道题）
- 第6次作业时间：5月27号~6月10号（1道题）
- 第2次lab时间：4月17日 – 6月15日
 - 基础任务（多个test programs通过率）
 - Bonus（占比Lab2的50%，简单AI卷积加速核设计，开放式、无强制要求）
- 期末汇报
 - 请自行挑选并阅读近5年内的芯片/架构等相关论文
 - 请同学们准备10页左右的PPT，每位同学12分钟演讲 + 3分钟答疑
 - 每位同学均为汇报同学打分，互评成绩占期末汇报成绩50%，请于5月19日前将汇报题目发给助教



目录

CONTENTS



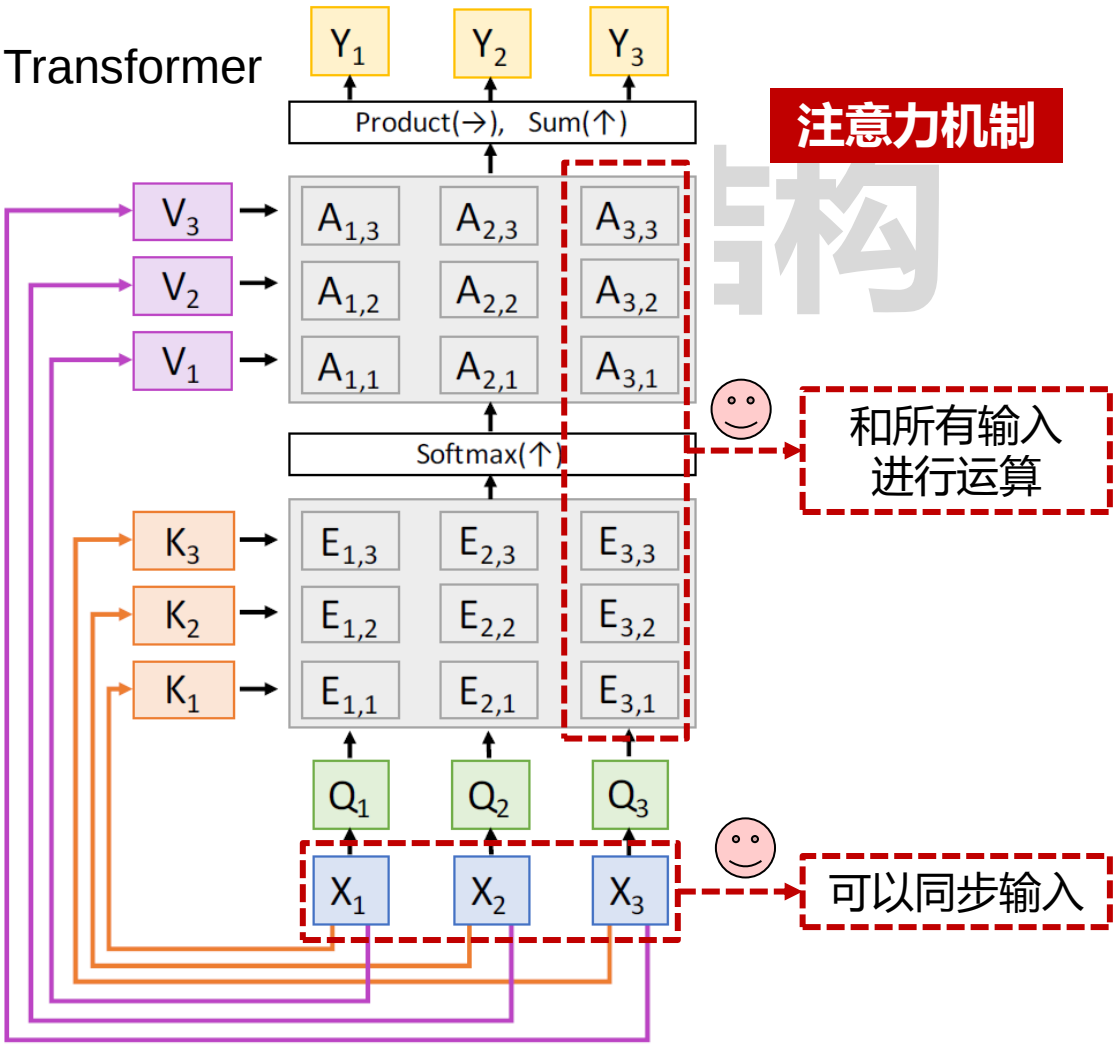
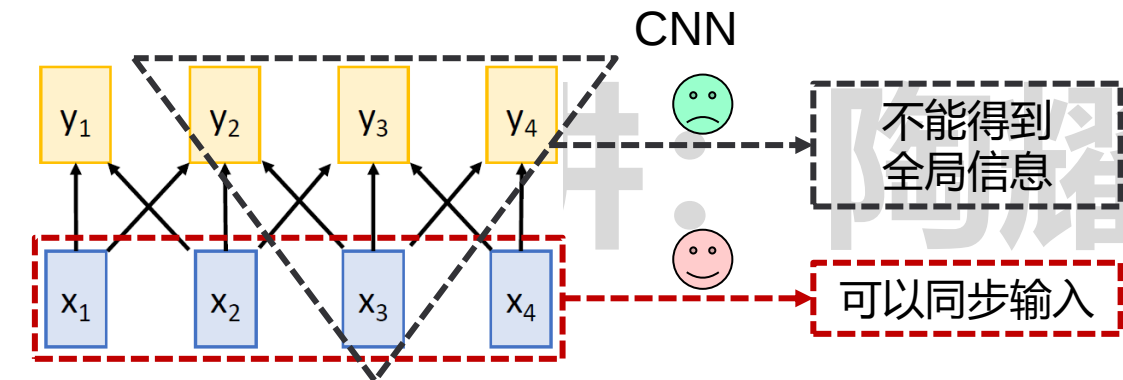
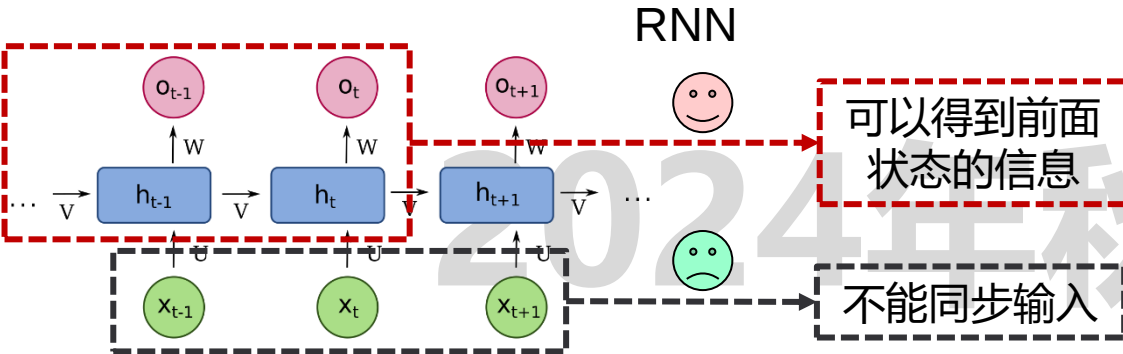
- 01. 典型AI芯片架构**
- 02. AI芯片软硬件协同设计**
- 03. AI大模型芯片设计**
- 04. 存算一体AI芯片架构**

研究背景

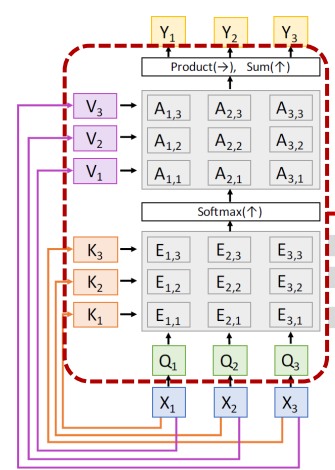


Transformer模型优势

指标	RNN	CNN	Transformer
数据特征	有序序列	多维网格	一组向量
长序列处理	✓	X	✓
可并行性	X	✓	✓



Transformer硬件部署问题



时间复杂度
 $O(n^2 \times d)$
n: 序列长度
d: 维度

序列长度(32K)增大

任务种类增多 → 字典维度增大

注意力模块
延时增加

训练延时随模
型增大增加

输入维度增大
导致能耗增加

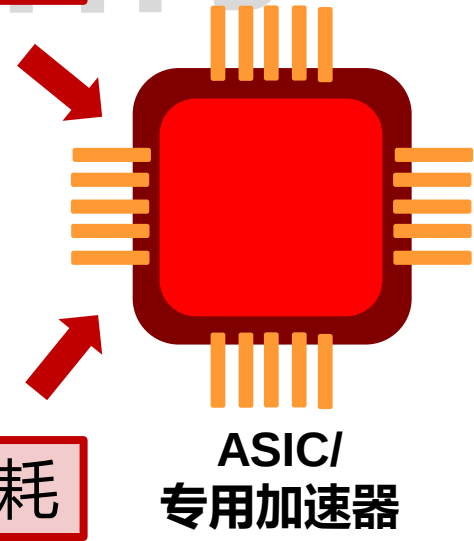
能耗随模型
增大增加

高延时

高能耗

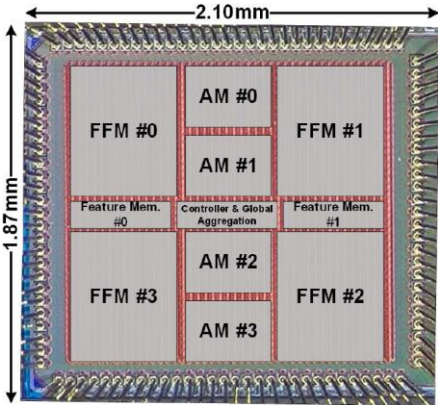
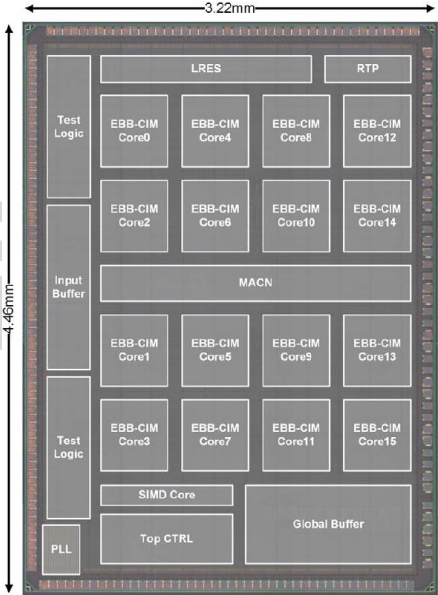
模型名称	训练时间(GPU-hours)	总能耗(MWh)	碳排放量(tCO ₂ eq)
OPT-175B	809472	365	137
BLOOM-175B	1082880	475	183
LLaMA	7B	82432	36
	13B	135168	59
	33B	530432	233
	65B	1022362	449
	70B	1720320	291
LLaMA-2	7B	184320	83
	13B	368640	162
	34B	1038336	396
	70B	1720320	291

From LLaMA/LLaMA 2, 数据收集于A100-80GB



Transformer专用加速器

工作	单位	存储介质	存算模式	加速模型	流片/仿真
ReTransformer	杜克大学	ReRAM	存内计算	Transformer	仿真
ISSCC'2023	复旦大学	SRAM	近存计算	Transformer	流片
ReBERT	KAIST	ReRAM	存内计算	BERT	仿真
TranCIM	清华大学	SRAM	近存计算	BERT	流片
X-Former	普渡大学	SRAM/RRAM	近存/存内计算	BERT	仿真
P3ViT	澳门大学	SRAM	存内计算	ViT	流片
H3DAtten	佐治亚理工	SRAM/RRAM	近存/存内计算	Swin Transformer	仿真
MuTCIM	清华大学	SRAM	近存计算	MML Transformer	流片
TransPIM	加利福尼亚大学	DRAM	近存计算	BERT+GPT	仿真



Liu, Shiwei, et al. ISSCC 2023

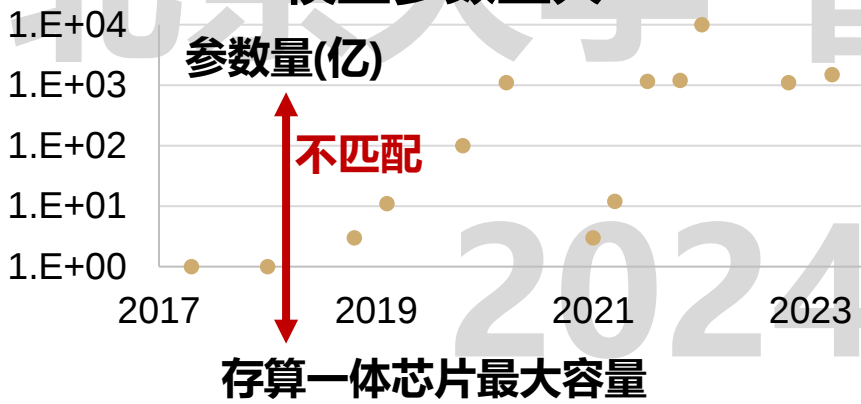
Tu, Fengbin, et al. ISSCC 2023

Transformer专用加速器挑战

高延时&高能耗 → ASIC/专用加速器

挑战1

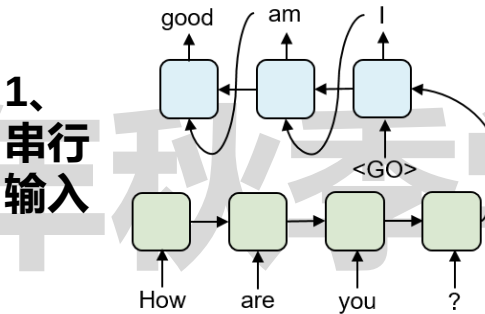
模型参数量大



模型级加速
减少计算量

挑战2

数据调度和传统深度神经网络不同



模块级加速
高效数据调度

挑战3

缺乏对专有算子的支持

主要线性算子	非线性算子
矩阵向量乘法(MVM)	GELU
矩阵矩阵乘法(MatMul)	LayerNorm
残差相加	Softmax
输入和权重不固定	动态归一化

算子级加速
加速专有算子

Transformer网络结构

□ 基本结构:

- 多个编码器层-多个解码器层

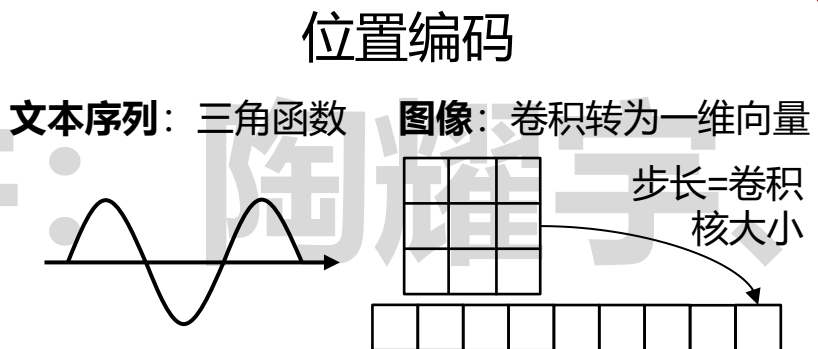
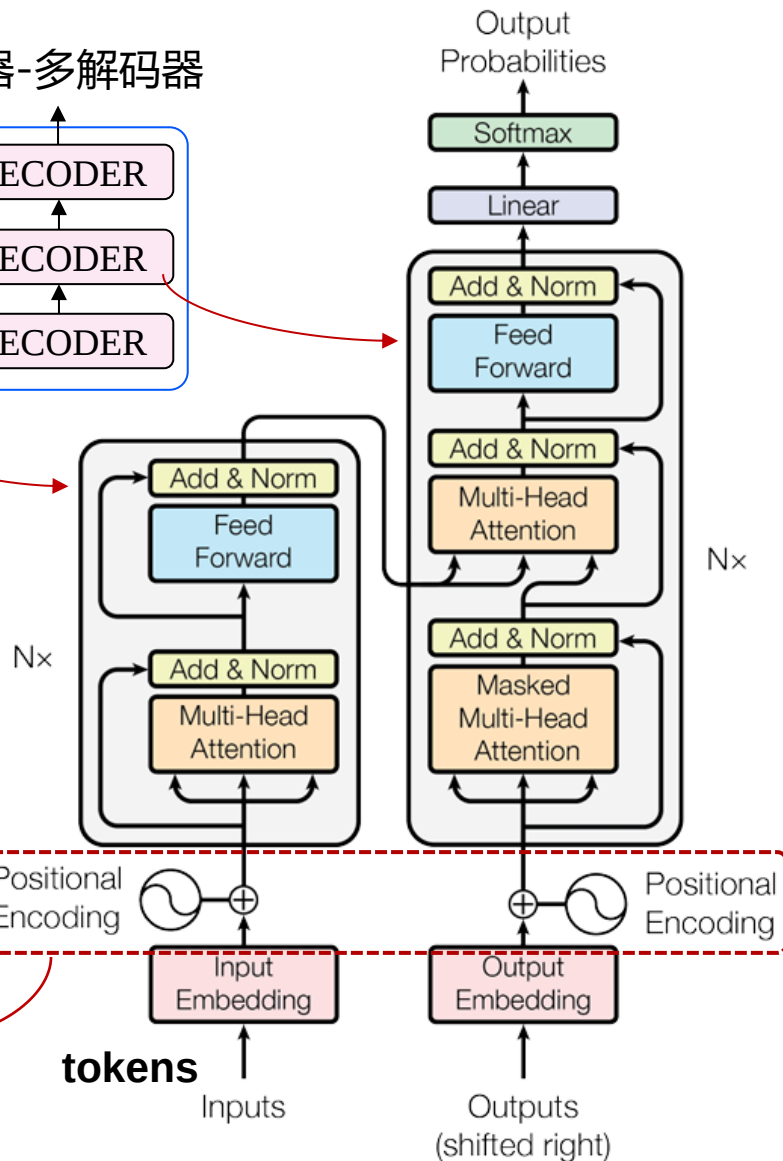
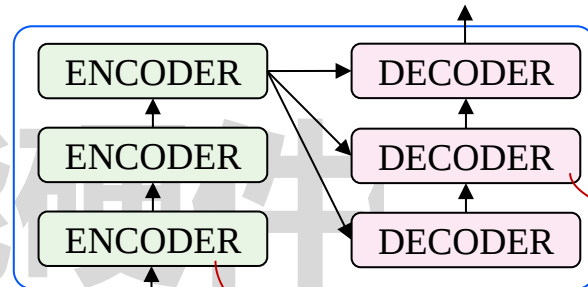
□ 输入序列:

- 由一组token组成
- 每个token为一个向量表示一个词根
- 序列开头[CLS]——句子向量表示
- 输入模型前经过embedding层

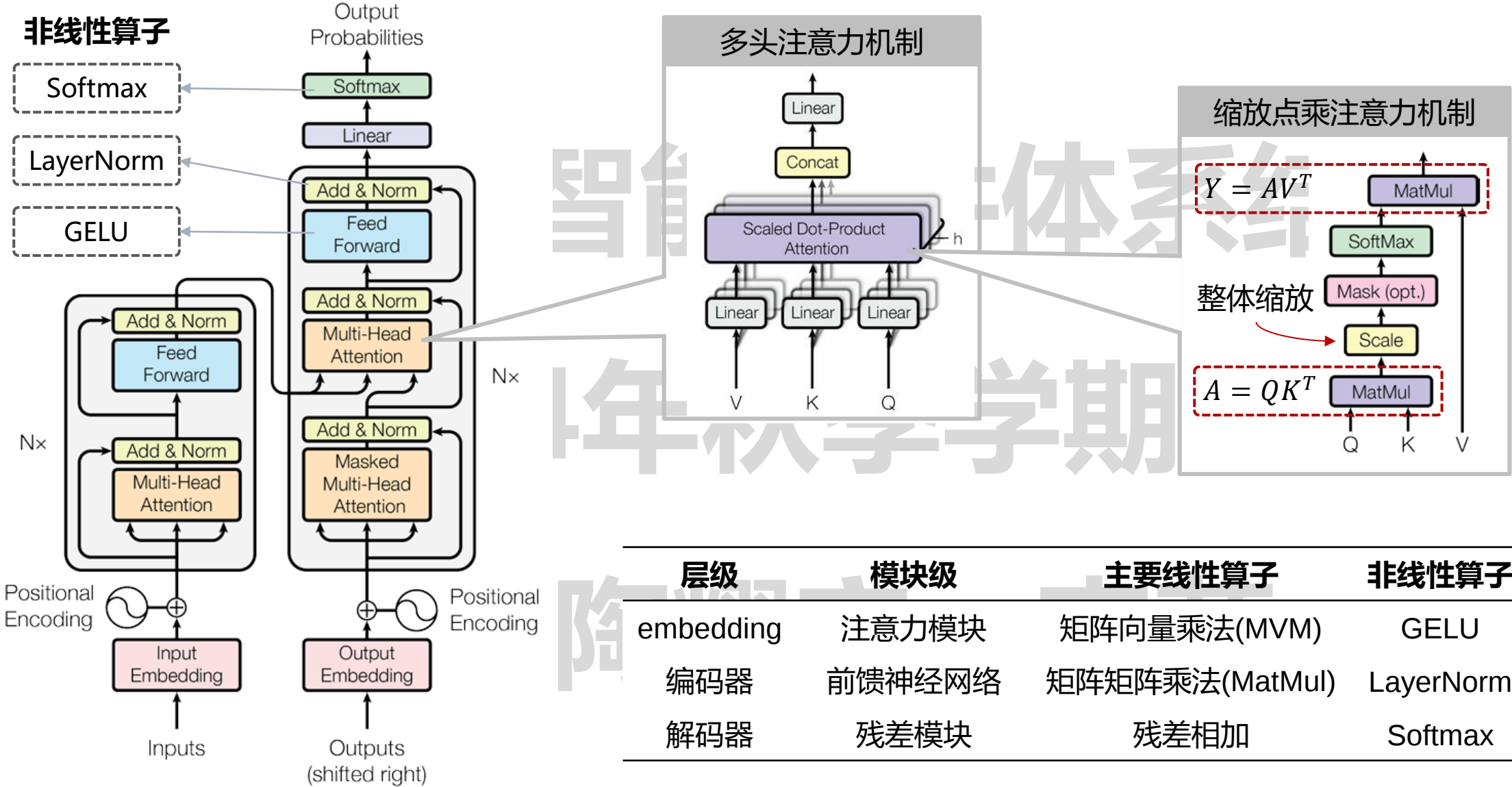
□ embedding:

- 输入编码: 字典信息
- 位置编码: 位置信息

基本结构: 多编码器-多解码器



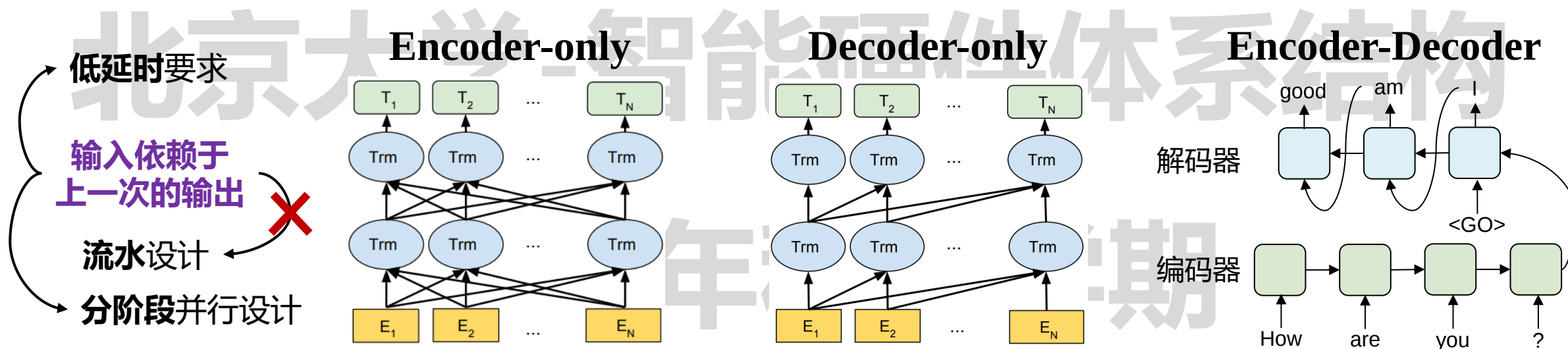
Transformer网络结构



Transformer种类

- 根据**网络结构**可分为：

- Encoder-only; Decoder-only; Encoder-Decoder



常见模型	BERT, RoBERTa	GPT, LLAMA	T5, BART
适合任务	分类任务	问答任务	翻译任务
输入输出特征	所有token并行输入	输入依赖于上一次的输出	输出强依赖于输入
适配加速器特征	高吞吐	低延时	低延时
适配并行设计	流水并行设计	分阶段并行设计	分阶段并行设计

目录

- 研究背景
- 模型级加速
 - 模型压缩
 - 稀疏性处理
- 模块级加速
 - 注意力模块加速
 - 前馈神经网络加速
- 算子级加速
 - 线性算子加速
 - 非线性算子加速
- 总结及展望

北京大学-智能硬件体系结构

2024年秋季学期

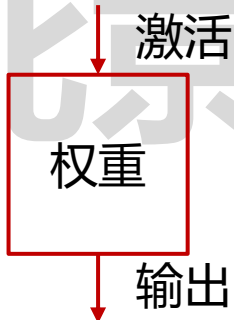
主讲：陶耀宇、李萌

模型级加速：模型压缩

• 模型压缩：硬件部署前算法优化

①

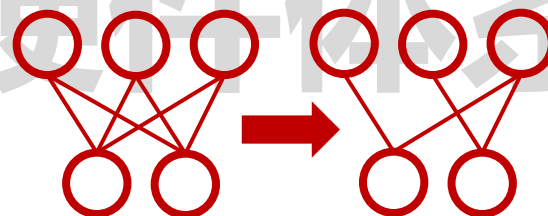
模型量化



- 对**激活/权重/输出**进行量化
- 降低模型**权重存储**需求
- 降低硬件**计算精度**要求

②

静态剪枝



- 删去**不重要的**注意力维度/头
- 降低硬件**计算量**

③

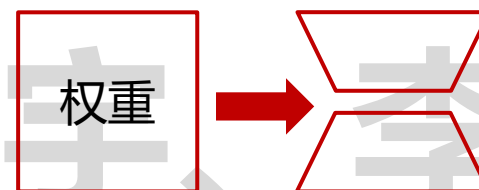
权重复用



- **重复利用**网络中的某一部分
- 降低模型**权重存储**需求
- 降低硬件**设计**需求

④

低秩分解



- 将权重矩阵拆分为几个低秩矩阵的组合
- 降低模型**训练存储**需求

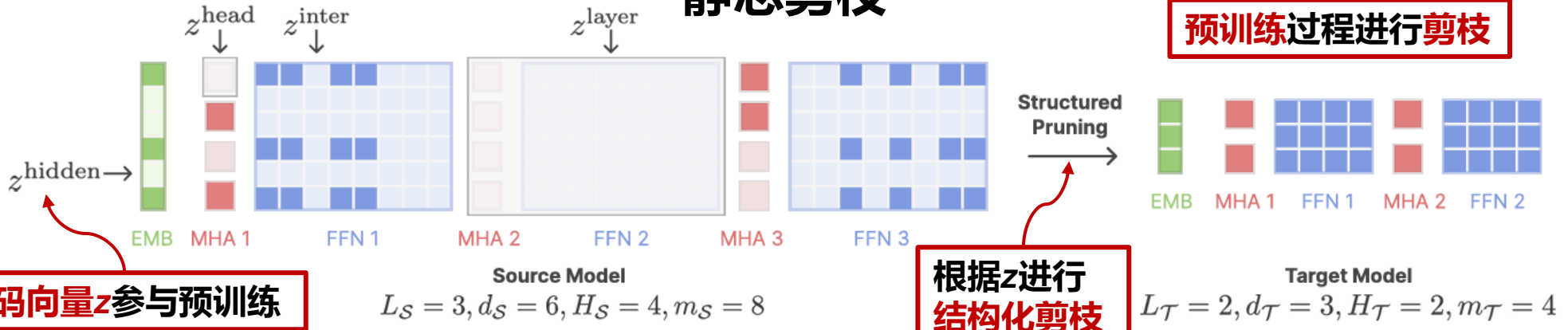
模型级加速：模型压缩

模型量化

Q-BERT	TernaryBERT	I-BERT	BinaryBERT	PTQ4ViT	FQ-ViT	BiT
INT8	三值	INT8	二值	INT8	INT4	二值

目前部分运算比如全连接层可实现三值/二值精度下降5%以内，INT8无损

静态剪枝



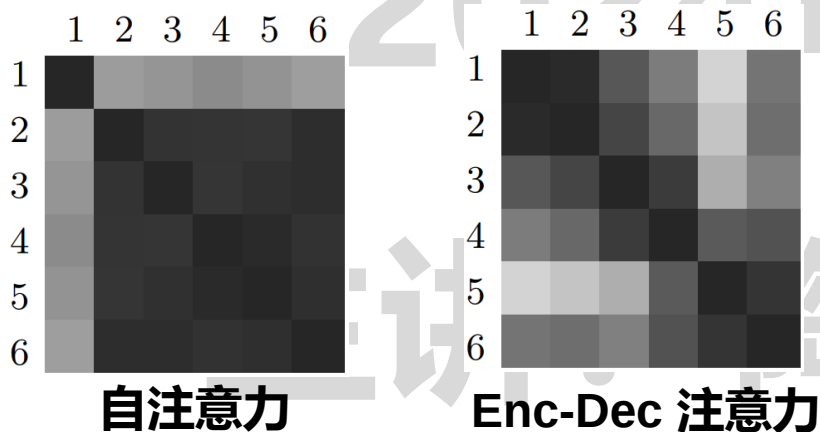
模型级加速：模型压缩

• 权重复用

注意力模块权重相似

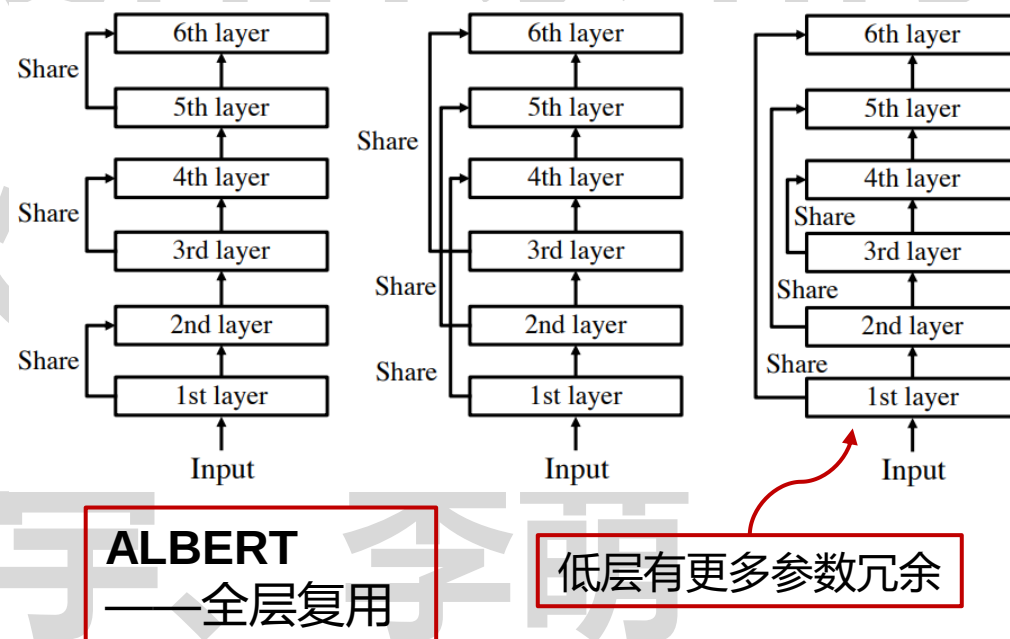
1、不同**维度**之间权重相似
多个维度学习同样的东西

2、不同**层**之间权重相似
不同层的权重的**JS**散度
颜色越深代表相似度越高



层间/维度间权重复用

根据**权重相似性**进行复用



模型级加速：模型压缩

• 低秩分解

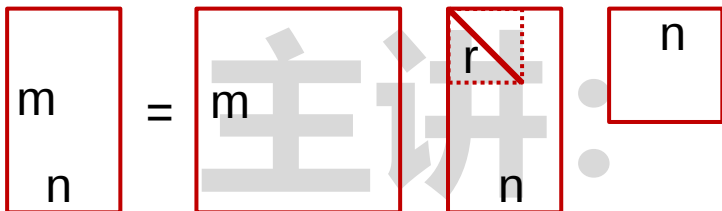
- 适用于**微调训练**过程
- 训练中的权重更新为**低秩**矩阵，可以用两个更小矩阵来表示
- 拆分更新权重降低**训练显存**

矩阵低秩分解

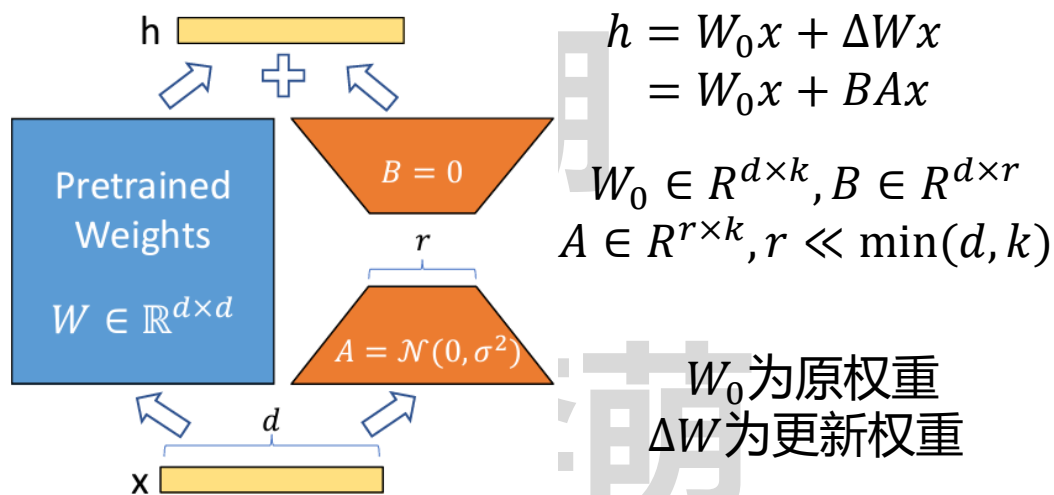
□ 奇异值分解SVD

$$A = U\Sigma V^T$$

仅有 r 宽度的对角
线上为特征值



微调训练应用



目录

- 研究背景
- 模型级加速
 - 模型压缩
 - 稀疏性处理
- 模块级加速
 - 注意力模块加速
 - 前馈神经网络加速
- 算子级加速
 - 线性算子加速
 - 非线性算子加速
- 总结及展望

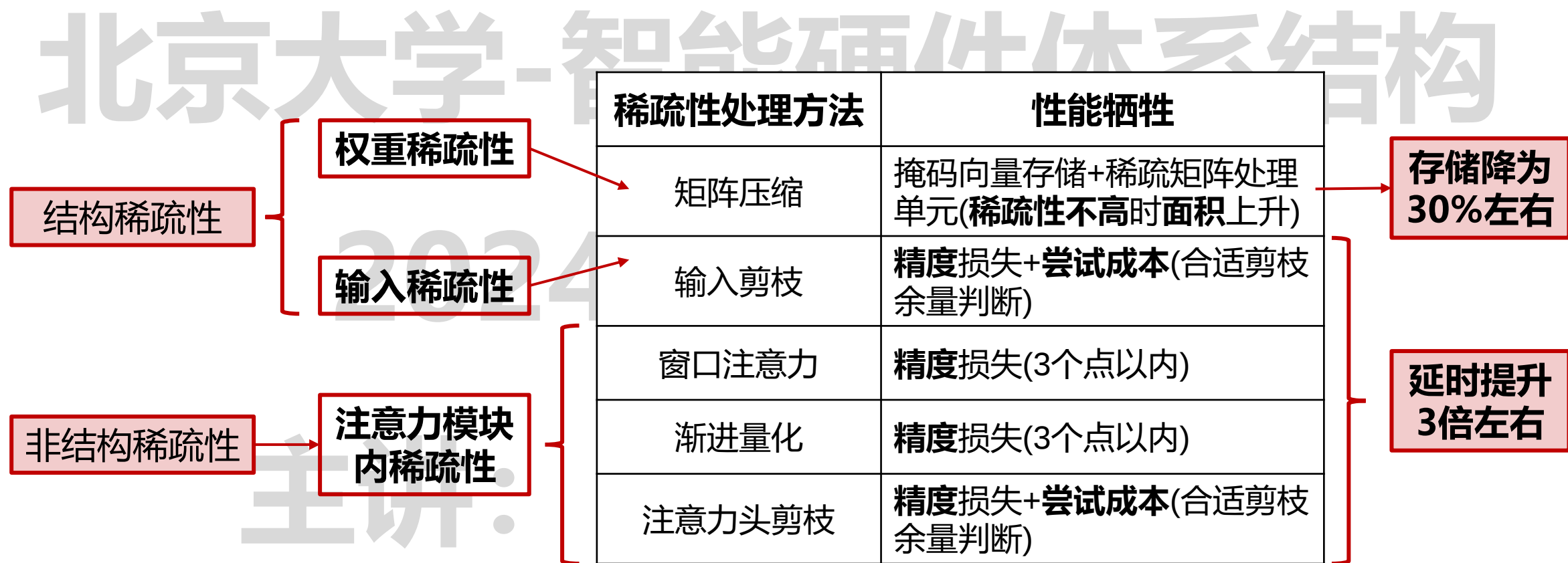
北京大学-智能硬件体系结构

2024年秋季学期

主讲：陶耀宇、李萌

模型级加速：稀疏性处理

- 稀疏性处理：硬件部署后**电路**优化
 - 针对**不同稀疏性**优化



模型级加速：稀疏性处理

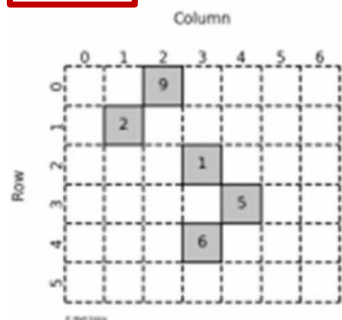
• 权重稀疏性

- 根据稀疏向量直接对输入进行裁剪

稀疏矩阵

- ❑ **CSC**: 按列压缩
- ❑ **CSR**: 按行压缩
- ❑ **COO**: 存储坐标

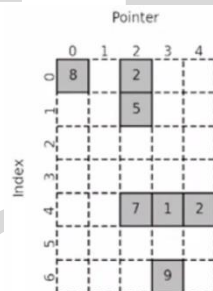
坐标



COO



CSC

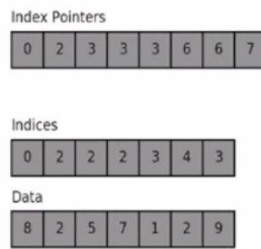
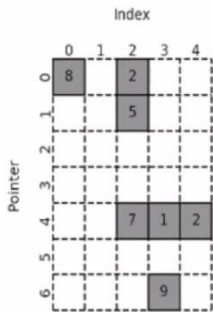


CSC



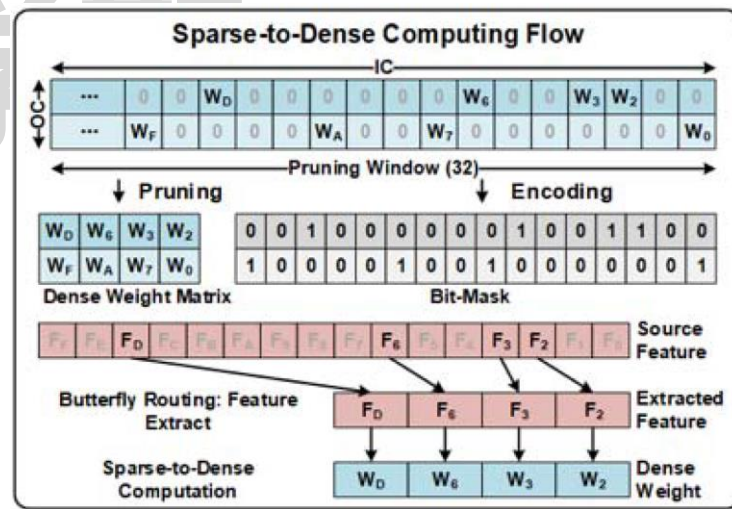
非零值的行的索引+行的起始位置指针

CSR



稀疏性矩阵应用

- ❑ 稀疏矩阵 $\rightarrow$ 压缩矩阵 + 掩码向量
- ❑ 根据掩码向量对输入进行裁剪



模型级加速：稀疏性处理

• 输入稀疏性

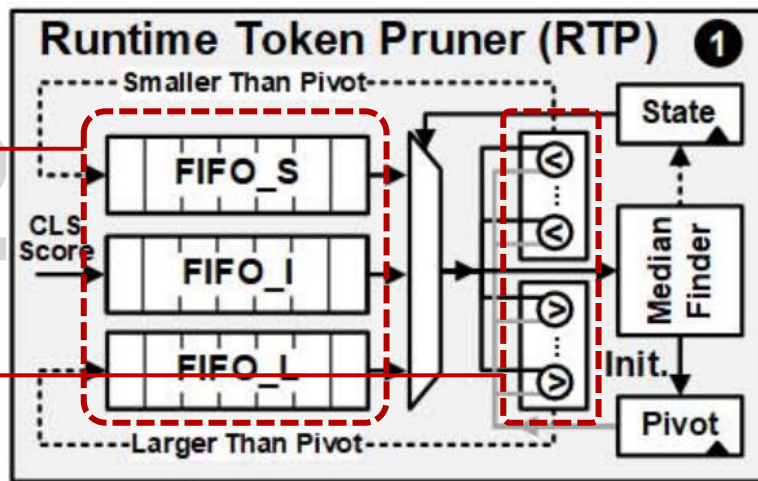
- 删去相对不重要的输入

三块FIFO

- FIFO_L: 大于阈值的token
- FIFO_I: 等于阈值的token
- FIFO_S: 小于阈值的token

一组比较器

- 比较对象固定为阈值
- 一组token进行并行比较



操作流程

- 循环实现
- 循环中止目标:
 - FIFO_S/FIFO_L内元素个数和设定值相当

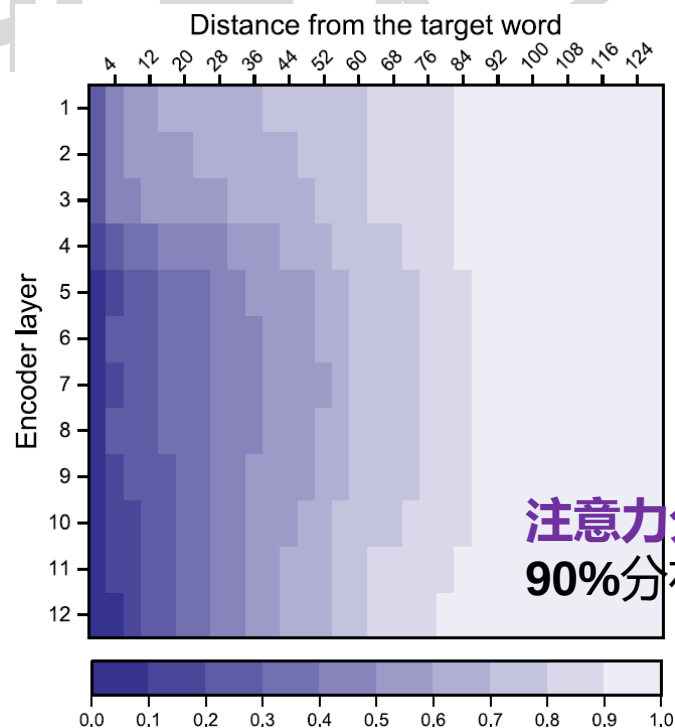
循环内容

- 初始化:
 - 阈值为[CLS]
 - 所有token放入FIFO_I
 - 选定FIFO_I
- 重复:
 - 阈值比较并放入FIFO
- 循环判断:
 - 满足要求则输出
 - 不满则选定FIFO_S/L
 - 阈值为其中随机一个token

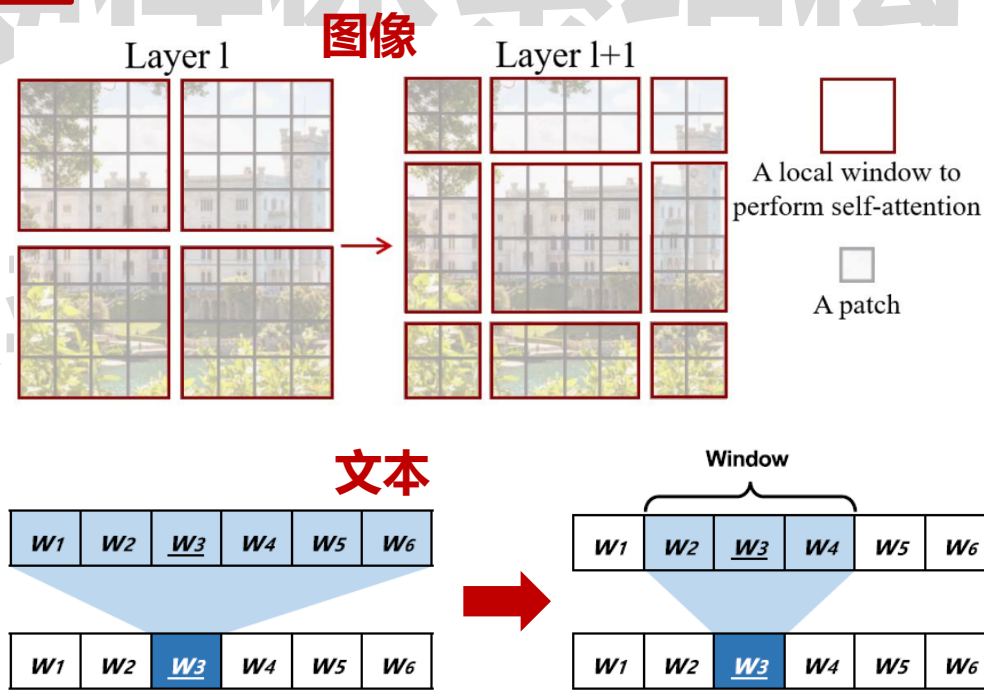
模型级加速：稀疏性处理

- 注意力模块内稀疏性
 - 利用**注意力机制**更关注于关键词的特征

窗口注意力



- 按注意力模块更关注目标词**附近**
- 目标词可以只和**附近token**进行注意力计算



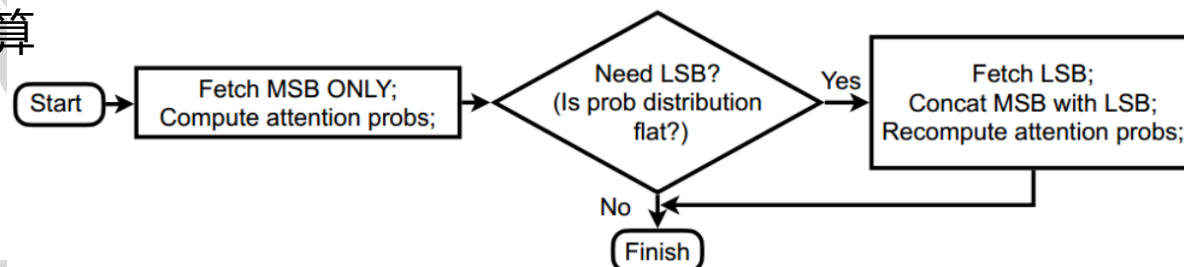
模型级加速：稀疏性处理

- 注意力模块内稀疏性

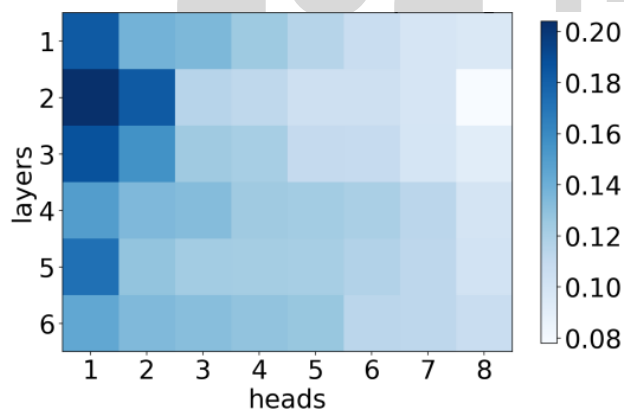
- 利用**注意力机制**更关注于关键词的特征

渐进量化

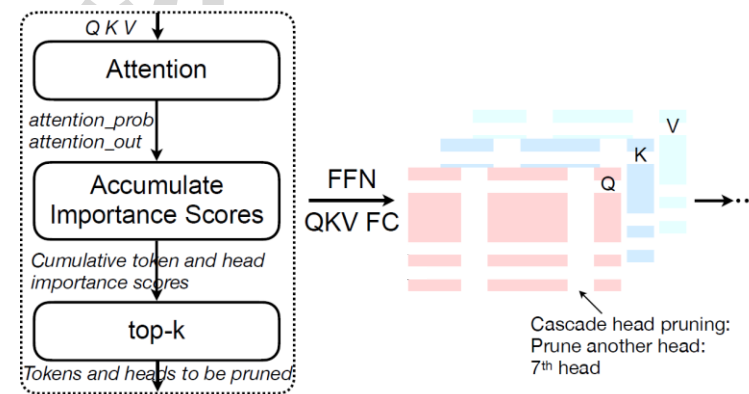
- 按**位**进行计算，而不是按**token**计算
- 每次计算**输入/权重**的一个比特
- 判断输出是否**平整**
- 不平整——已经实现注意力功能



注意力头剪枝



- 各层各头注意力分数累加
- 各头累加结果**相差很大**
- 大小关系在**各层统一**
- 提前进行注意力头剪枝



目录

- 研究背景
- 模型级加速
 - 模型压缩
 - 稀疏性处理
- 模块级加速
 - 注意力模块加速
 - 前馈神经网络加速
- 算子级加速
 - 线性算子加速
 - 非线性算子加速
- 总结及展望

北京大学-智能硬件体系结构

2024年秋季学期

主讲：陶耀宇、李萌

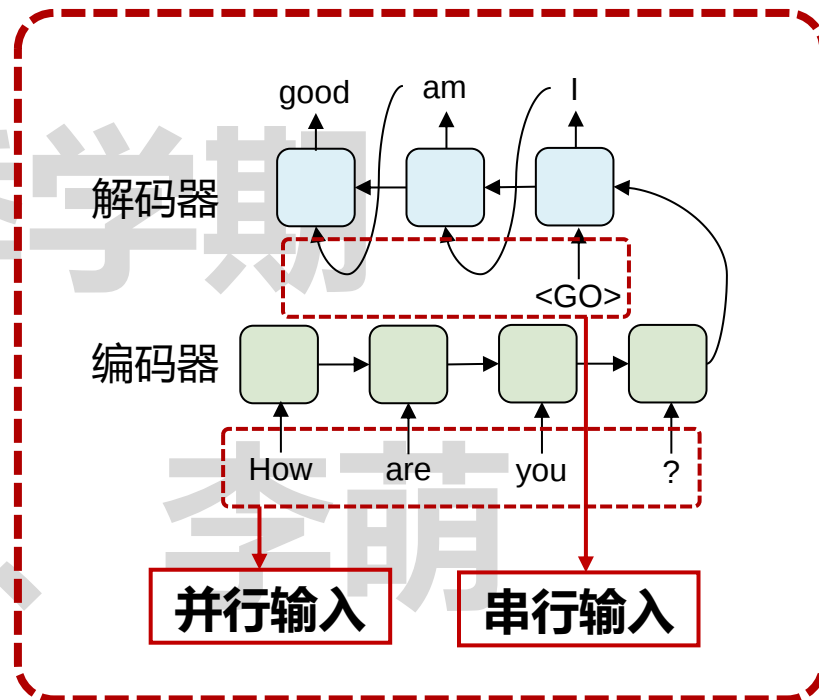
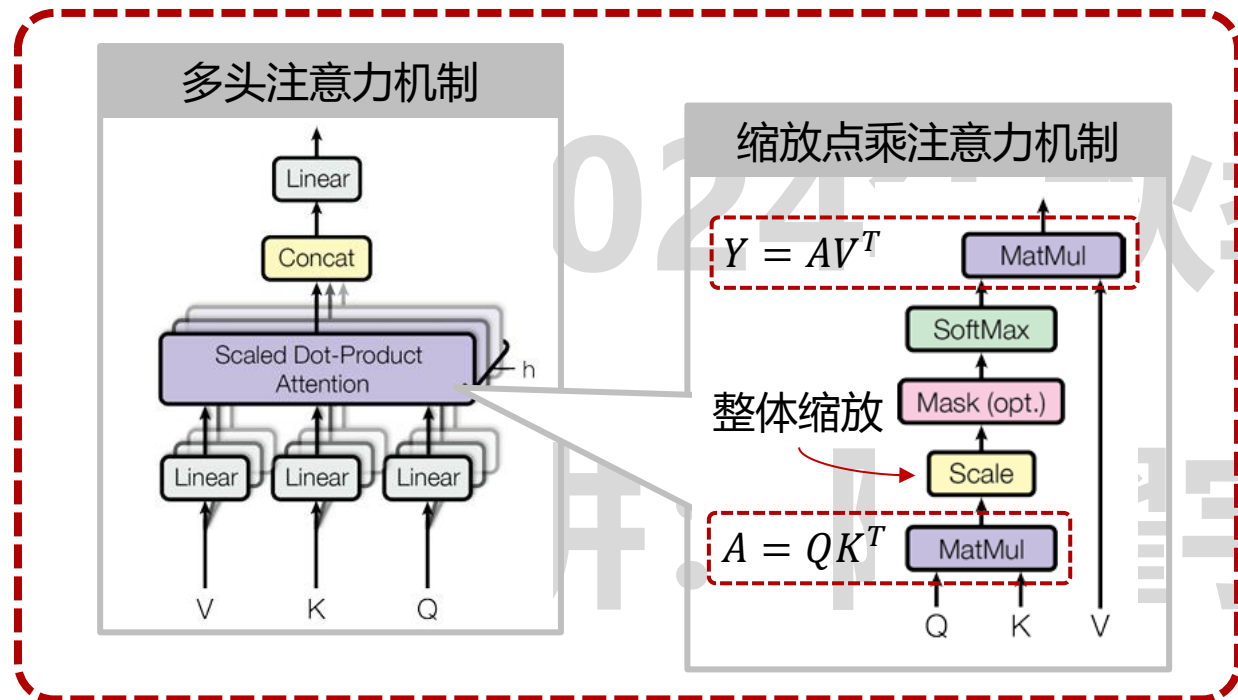
模块级加速：注意力模块加速

• 注意力模块结构分析

- 数据调度主要难点：MatMul运算需要等待所有token的线性层运算结束

• 架构加速需求：

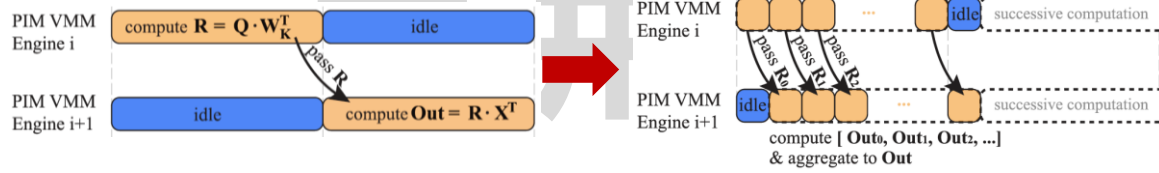
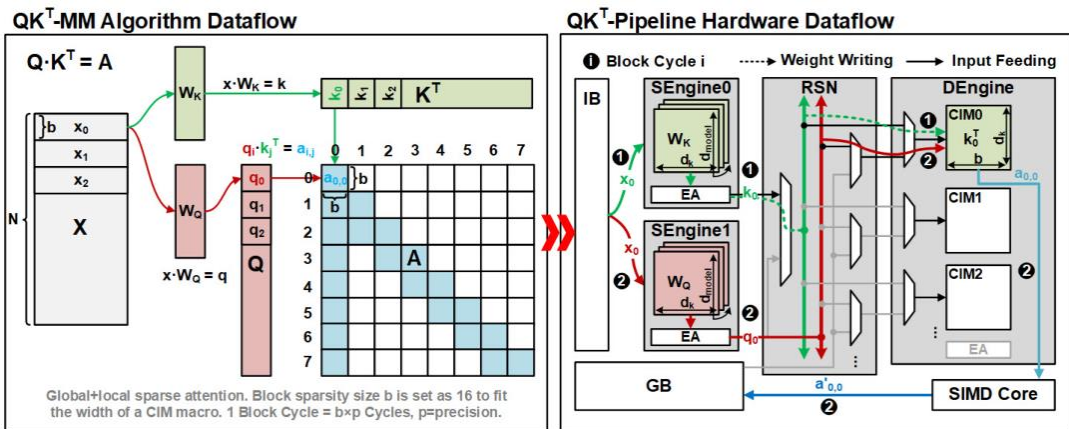
- 流水并行设计、分阶段并行设计



模块级加速：注意力模块加速

- 流水并行设计
 - 将输入拆分为**多块切片**提高计算核心利用率
 - Q/K/V线性层**运算错位**以并行进行写入权重和计算

输入切片



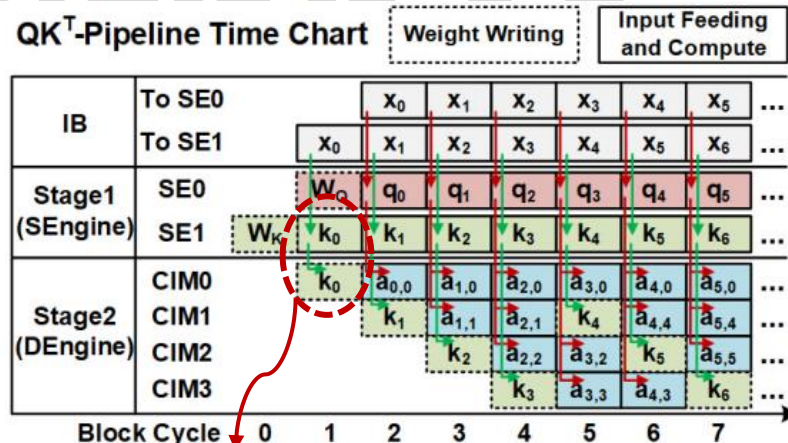
流水设计



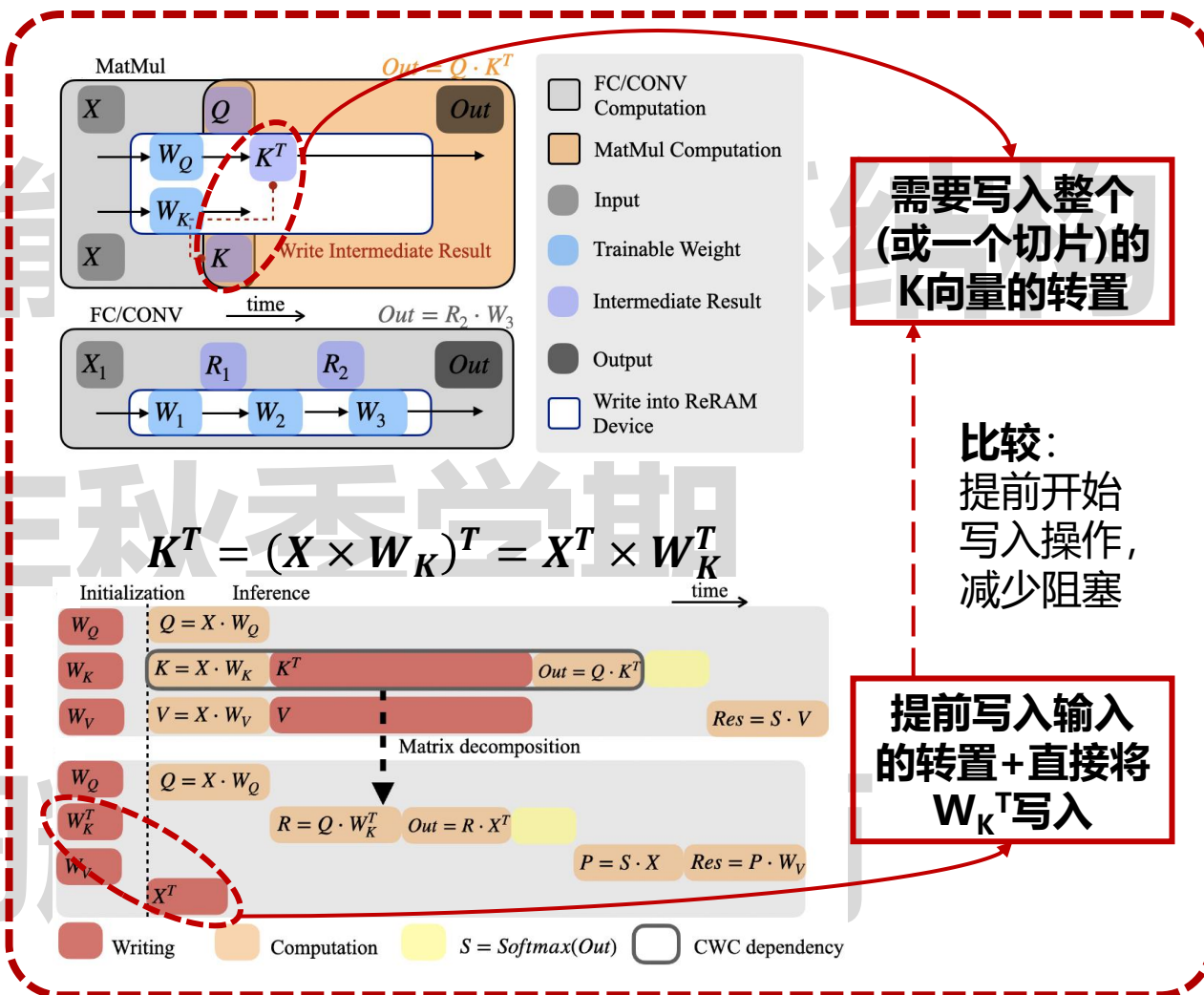
模块级加速：注意力模块加速

• 流水并行设计

- 通过拆分运算以提前写入

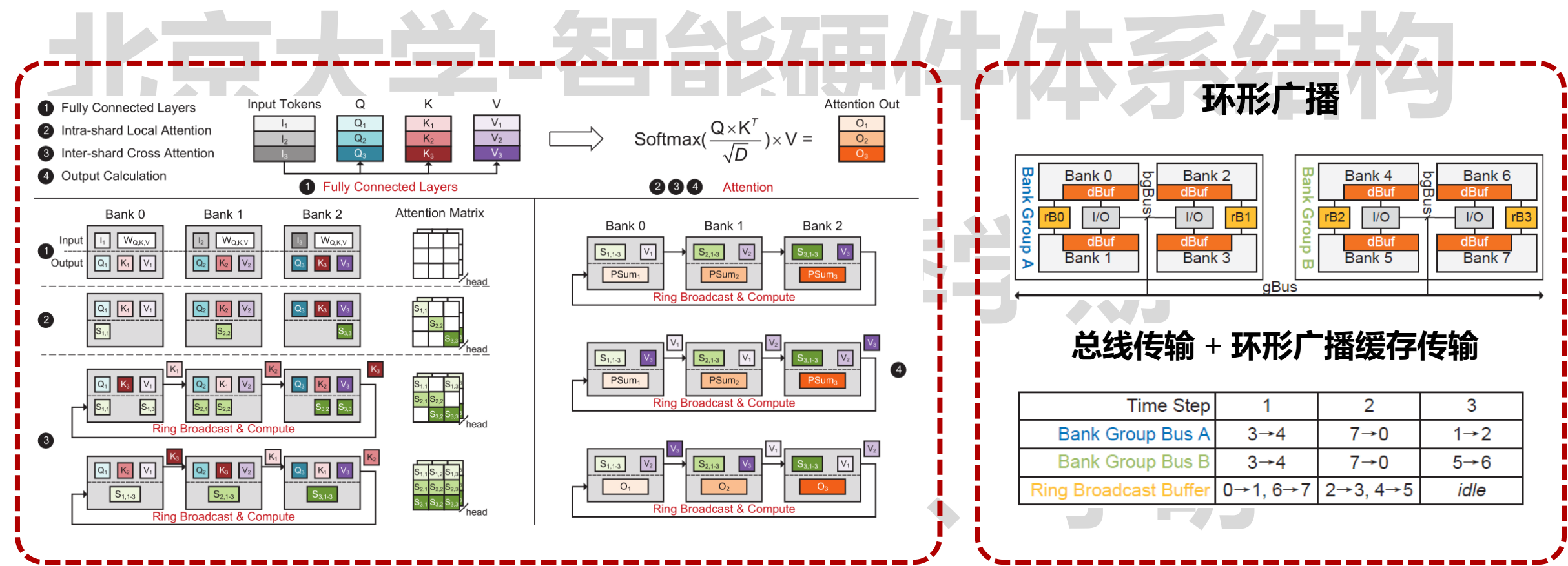


写入时间很短



模块级加速：注意力模块加速

- 分阶段并行设计
 - 通过**环形广播**和所有token进行注意力计算



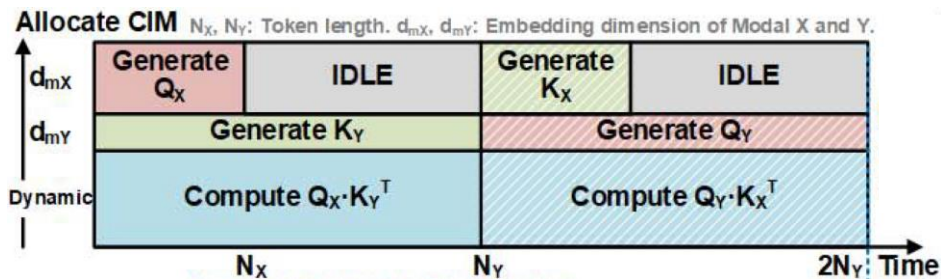
模块级加速：注意力模块加速

分阶段并行设计

- 针对**多模态**输入规划权重片上分布
- 提高**计算核心利用率**，降低推断延时

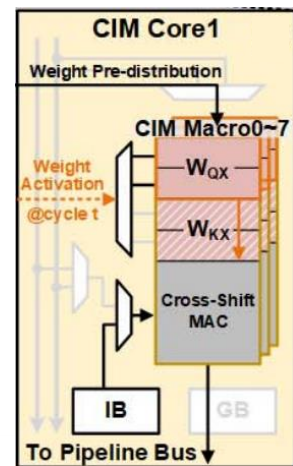
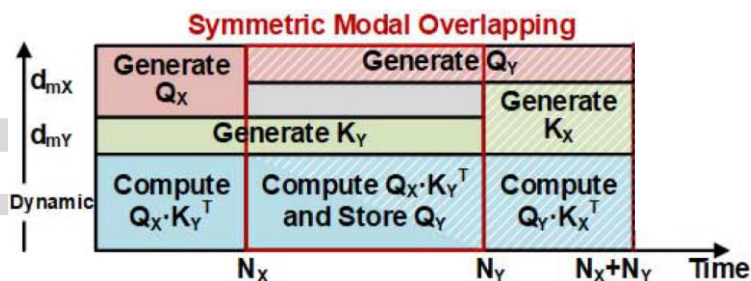
多模态输入的问题

- 传统权重分布：不同步骤在不同计算核心实现
- 多模态输入：不同阶段运算延时不同
- 导致多模态输入不匹配时产生**额外延时**



解决方案

- 计算核心**存储多份权重矩阵**
- 利用**模态对称性**重叠运算
- 分配负载提高硬件利用率



目录

- 研究背景
- 模型级加速
 - 模型压缩
 - 稀疏性处理
- 模块级加速
 - 注意力模块加速
 - 前馈神经网络加速
- 算子级加速
 - 线性算子加速
 - 非线性算子加速
- 总结及展望

北京大学-智能硬件体系结构

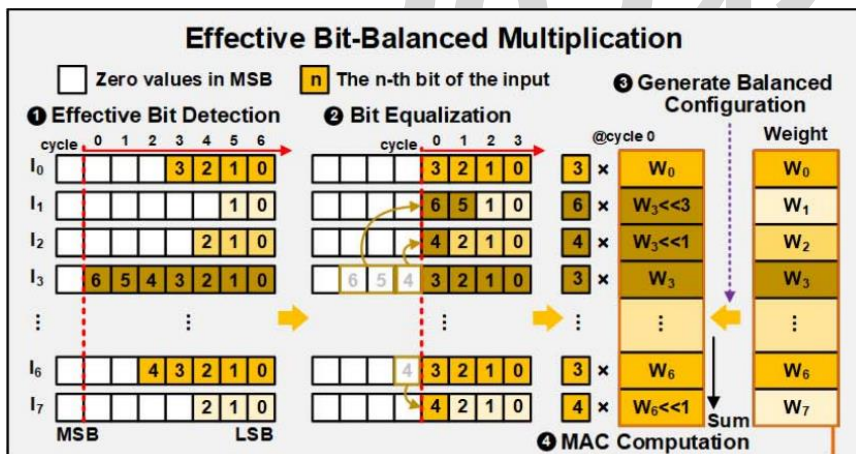
2024年秋季学期

主讲：陶耀宇、李萌

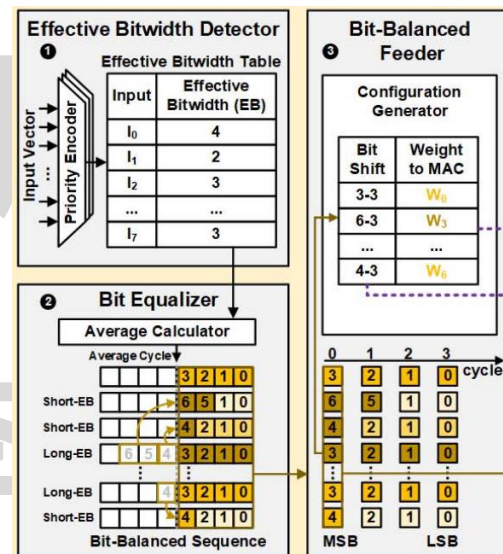
模块级加速：前馈神经网络加速

- 输入特征——token之间的值差距较大
 - 注意力机制本身设计目的是使模型更关注关键位置
 - softmax函数会放大较大输出，负值等较小输出会接近于0
- 输入前进行**有效位宽均衡**操作

有效位宽均衡



- 输入检测**有效位宽**
- 均衡有效位宽
- **等有效位宽序列**输入
- 根据**均衡表格**配置权重
- 根据均衡表格接收输出



目录

- 研究背景
- 模型级加速
 - 模型压缩
 - 稀疏性处理
- 模块级加速
 - 注意力模块加速
 - 前馈神经网络加速
- 算子级加速
 - 线性算子加速
 - 非线性算子加速
- 总结及展望

北京大学-智能硬件体系结构

2024年秋季学期

主讲：陶耀宇、李萌

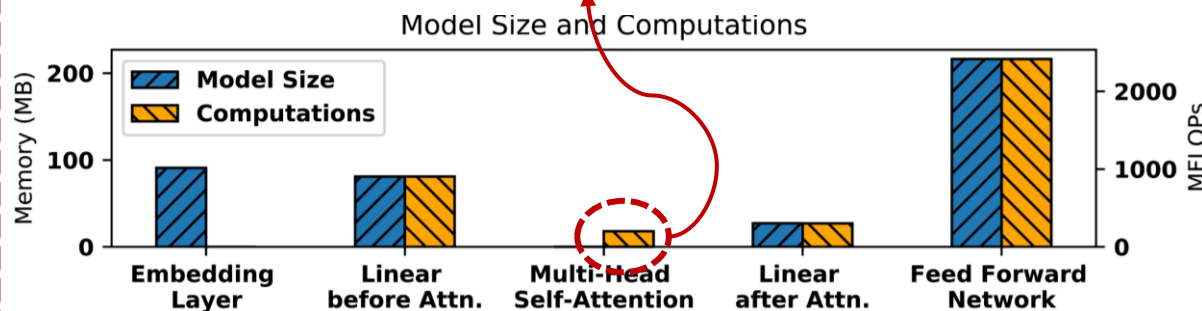
算子级加速：线性算子加速

• 线性算子加速

- 全连接层、注意力机制、ViT的卷积运算、残差等涉及常规乘加操作
- 以**MVM**算子和**MatMul**算子为主
- MVM算子：**权重矩阵固定**，输入向量可变
- MatMul算子：两个输入矩阵都不固定，每次计算前需要**加载权重矩阵**

算子计算量分布

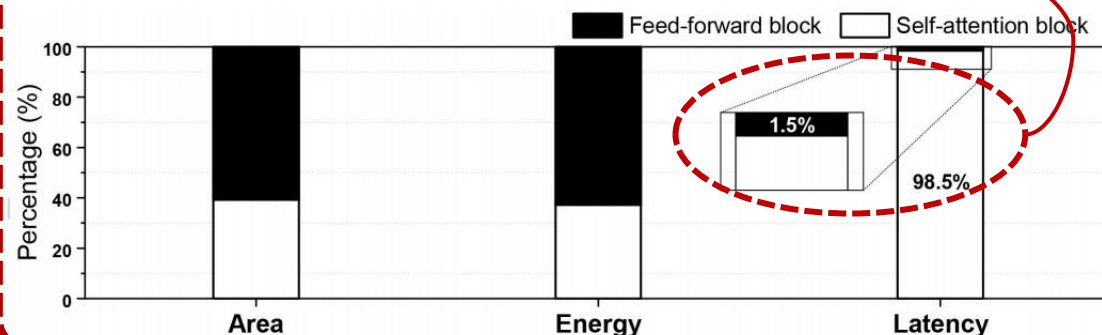
MatMul算子仅占0.5%左右



Ganesh, Prakhar, et al. *TACL* 2021

算子计算量分布

矩阵加载可能占据很高的延时



Kang, Myeonggu, Hyein Shin, and Lee-Sup Kim. *TCAD* 2021

算子级加速：线性算子加速

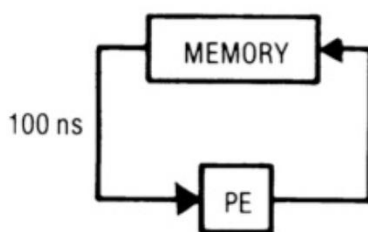
- 存算分离

- 脉动阵列加速乘法累加运算

脉动阵列概念

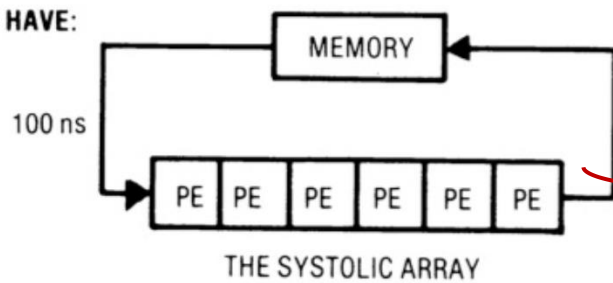
减少存储器调度时间

INSTEAD OF:



5 MILLION
OPERATIONS
PER SECOND
AT MOST

WE HAVE:



30 MOPS
POSSIBLE

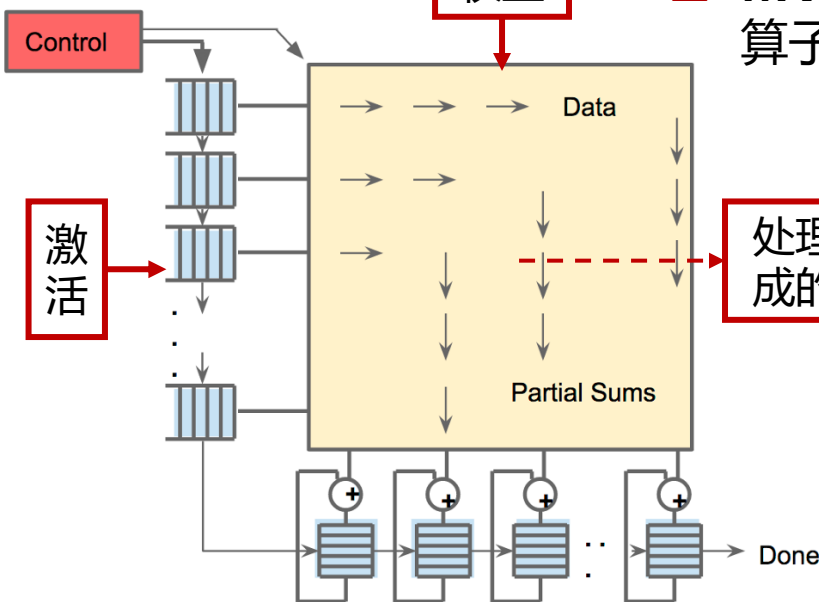
THE SYSTOLIC ARRAY

Kung, Hsiang-Tsung. *Computer* 1982

脉动阵列应用

权重

□ MVM/Matmul
算子都适配



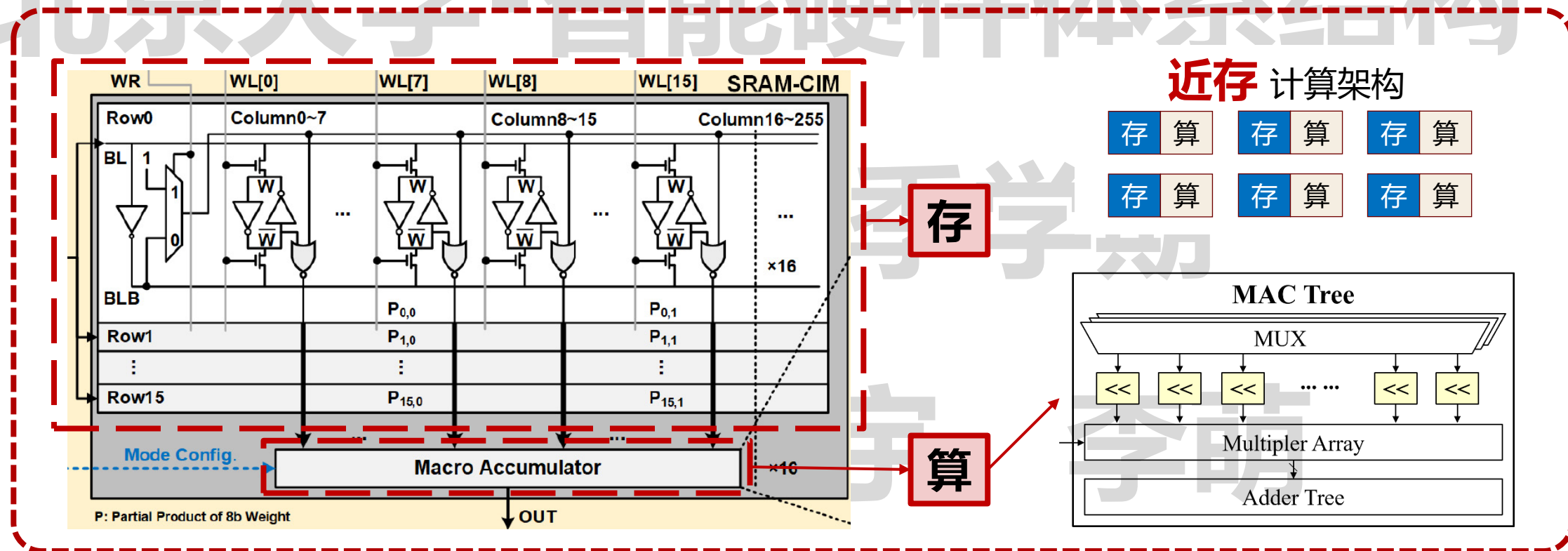
处理器单元构
成的二维网格

Jouppi, Norman P., et al. *ISCA* 2017

算子级加速：线性算子加速

• 近存计算

- 存储阵列输出处增加逻辑单元，**减少传输**；DRAM、SRAM等存储介质
- 乘法累加树——写入时间较长时影响MatMul算子实现

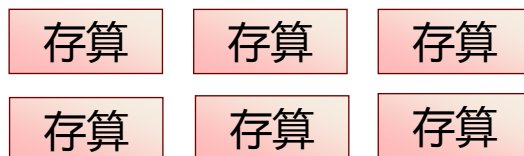
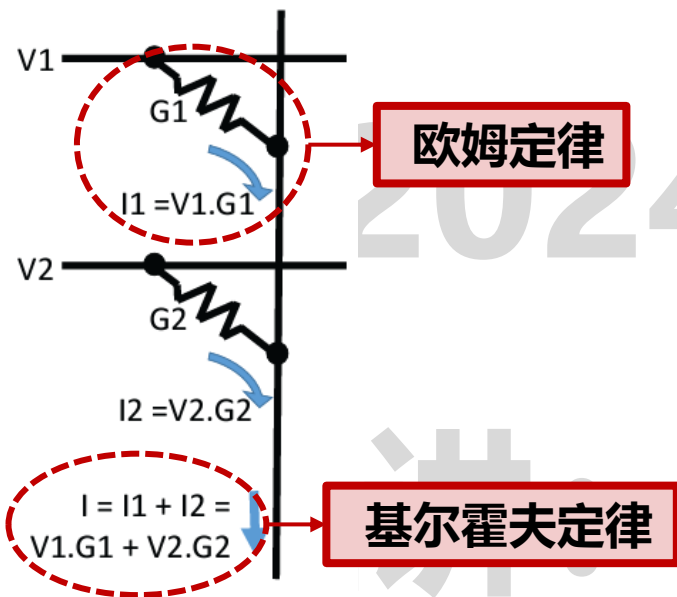


算子级加速：线性算子加速

• 存内计算

- 存储阵列内完成乘法累加运算，提高**吞吐**和**能效**；忆阻器、SRAM等存储介质
- 利用交叉阵列实现——写入时间较长时影响MatMul算子实现

存内 计算架构



- ❑ ① 模拟存内计算：
 - 开启多行直接实现累加计算
- ❑ ② 数字存内计算：
 - 单行开启，在阵列内完成乘法运算
- ❑ 算子实现：
 - MVM算子实现能效较高
 - 部分介质不适合MatMul算子



目录

- 研究背景
- 模型级加速
 - 模型压缩
 - 稀疏性处理
- 模块级加速
 - 注意力模块加速
 - 前馈神经网络加速
- 算子级加速
 - 线性算子加速
 - 非线性算子加速
- 总结及展望

北京大学-智能硬件体系结构

2024年秋季学期

主讲：陶耀宇、李萌

算子级加速：非线性算子加速

- 非线性算子加速

- 分为需要/不需要统计输入特征的部分

层级	模块级	主要线性算子	非线性算子
embedding	注意力模块	矩阵向量乘法(MVM)	GELU
编码器	前馈神经网络	矩阵矩阵乘法(MatMul)	LayerNorm
解码器	残差模块	残差相加	Softmax

输入特征统计加速

非线性函数加速

GELU

不需要统计输入特征

$$GELU(x) = 0.5x \left(1 + \operatorname{erf} \left(\frac{x}{\sqrt{2}} \right) \right)$$
$$\operatorname{erf}(x) = \frac{2}{\sqrt{\pi}} \int_0^x \exp(-t^2) dt$$

LayerNorm

需要统计输入特征

$$\operatorname{LayerNorm}(x) = \frac{x - \mu}{\sigma}$$
$$\mu = \frac{1}{C} \sum_{i=1}^C x_i, \sigma = \sqrt{\frac{1}{C} \sum_{i=1}^C (x_i - \mu)^2}$$

Softmax

需要统计输入特征

$$\operatorname{Softmax}(x) = \frac{\exp(x_i)}{\sum_{j=1}^K \exp(x_j)}$$
$$x = [x_1, \dots, x_k]$$

算子级加速：非线性算子加速

• 非线性算子加速：输入特征统计

- 均值统计/方差统计/指数统计

□ 均值统计

- 记录累和值和数量
- 每次数据流经过时更新

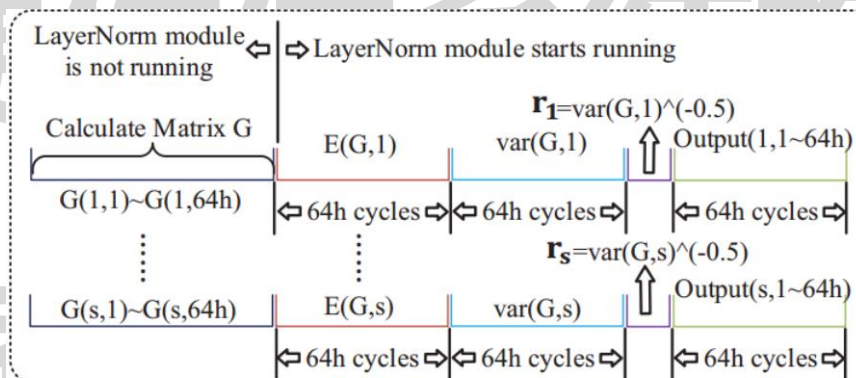
□ 方差统计

- 数学重构为 $var = \mu^2 - \frac{1}{C} \sum_{i=1}^C x_i^2$
- 记录平方值和数量
- 每次数据流经过时更新

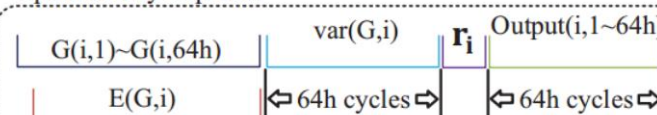
□ 指数统计

- 需要记录所有数据的指数
- 在数据输出时完成

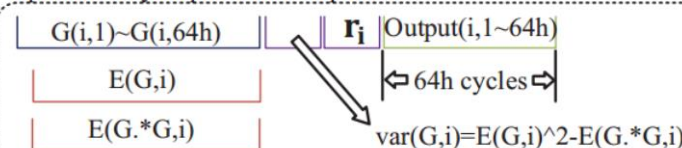
数据流上
进行加速



Optimized by step one:



Optimized by step one and step two:



算子级加速：非线性算子加速

• 非线性算子加速：非线性函数加速

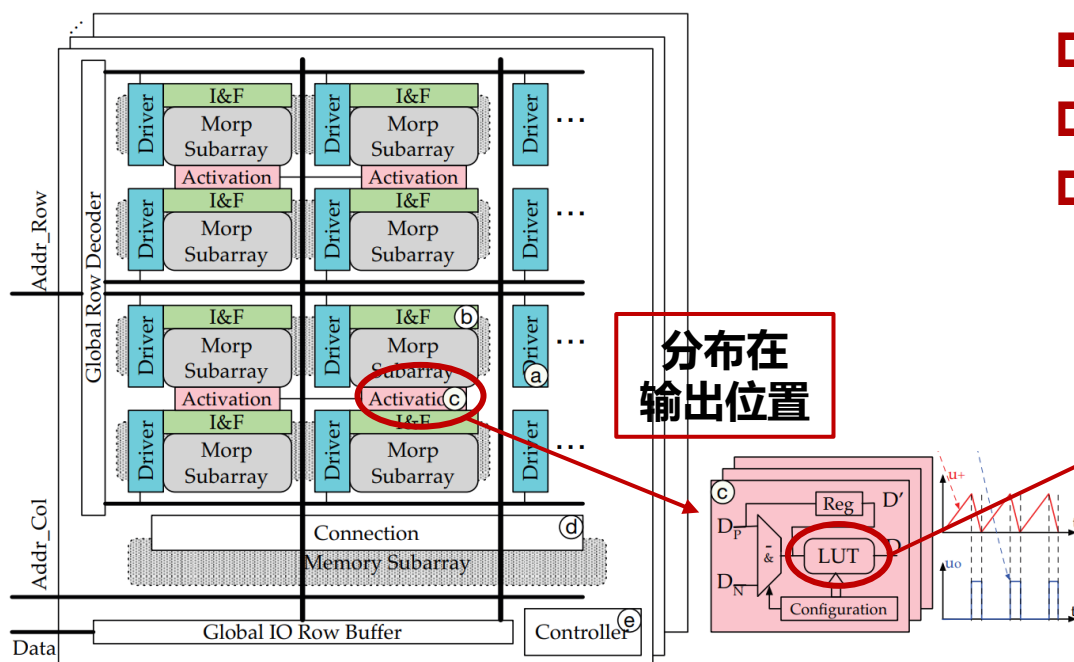
- 两种思路：将非线性函数变为线性函数/直接实现非线性函数



算子级加速：非线性算子加速

• 非线性函数实现：LUT(查找表)

- 优点：同样的电路结构可以用于实现不同的非线性函数，具有一定的灵活性
- 缺点：大范围输入需要大容量存储阵列



分布在
输出位置

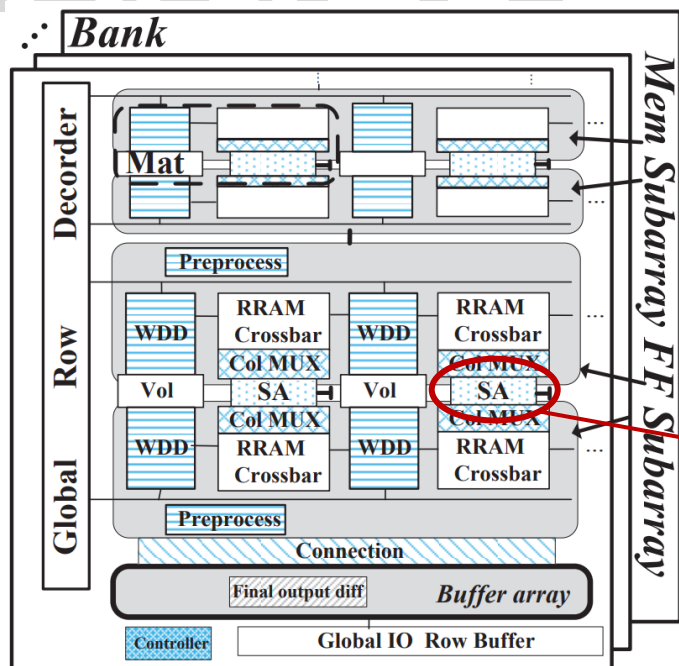
- 通过提前存储结果加速计算
- 可扩展性限制：大范围输入需要大容量
- 阵列大小限制：存储阵列大小限制计算范围
 - 拆分多个子阵列会导致高延迟



算子级加速：非线性算子加速

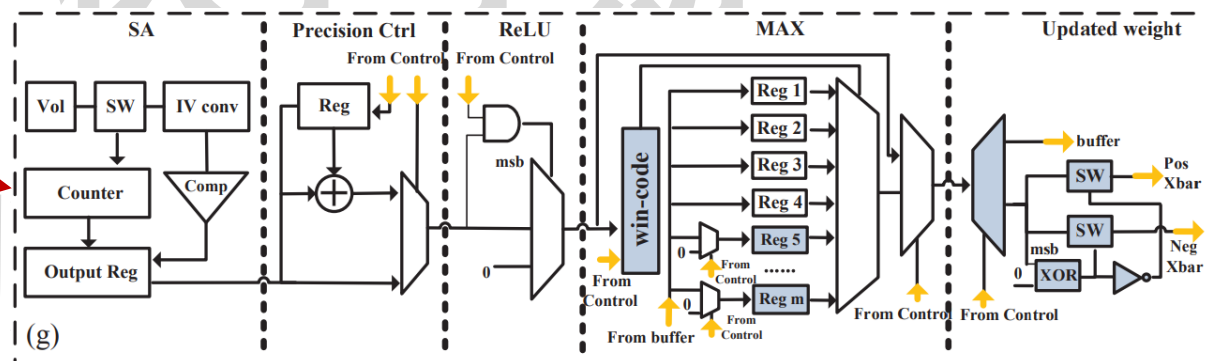
• 非线性函数实现：DLC(数字逻辑电路实现)

- 优点：计算过程**延迟**低+使用成熟CMOS工艺对**面积**要求较低
- 缺点：固化数字电路导致**灵活度**低+数字设计需要考虑**静态功耗**



分布在
输出位置

- 固化了数字逻辑电路
- 定制化：对面积要求低并且计算延时低
- 不定制化：通用计算核心或者泰勒级数实现
 - 前者电路复杂；后者引入乘法器，增加了面积

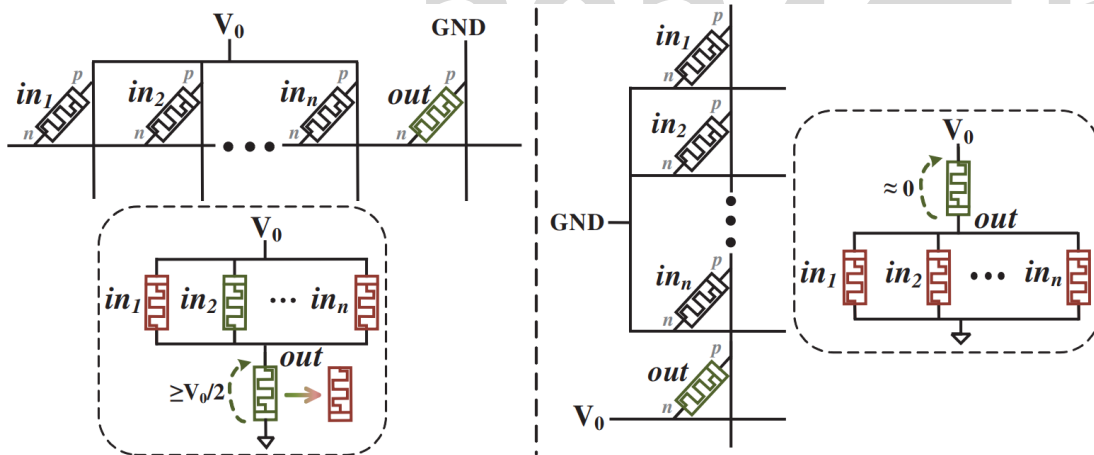


算子级加速：非线性算子加速

• 非线性函数实现：MLS&OO

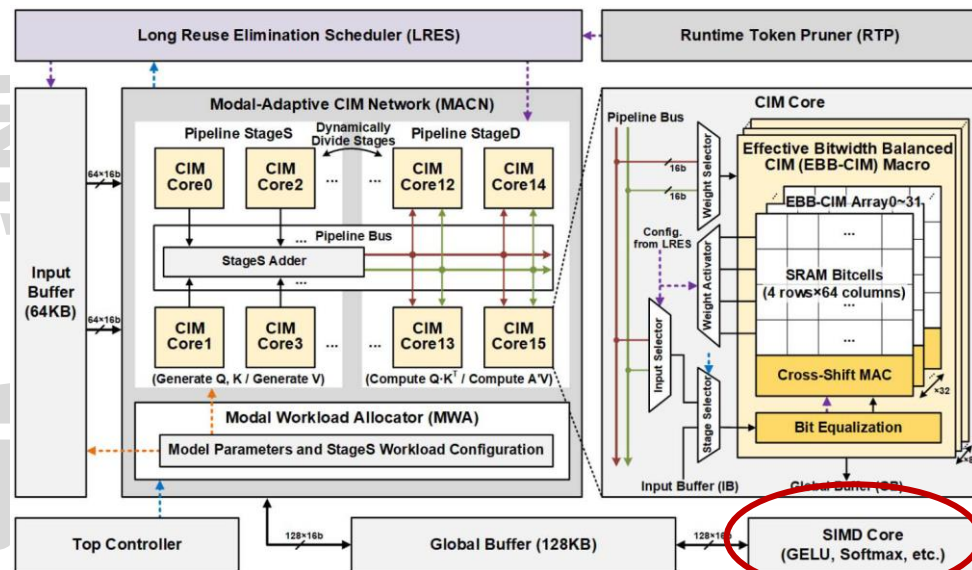
MLS(新器件实现逻辑)

- 通过器件分压执行写入操作
- 优点：静态功耗低+计算单元面积小
- 缺点：写入能耗高+定制化灵活度低+需要很多器件单元成本高



OO(操作片外卸载)

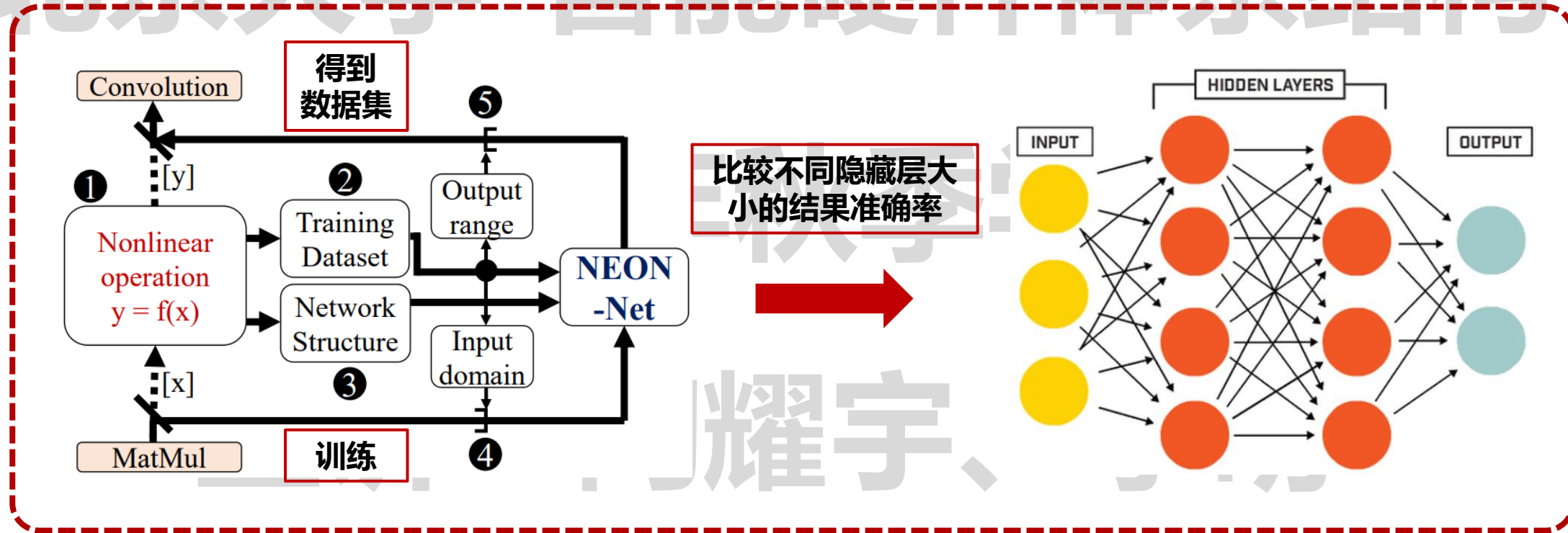
- 数据传输到片外完成
- 优点：直接利用通用CPU
- 缺点：数据移动过多影响能耗和延时



算子级加速：非线性算子加速

• 非线性函数转换：神经网络实现

- 利用存算一体等对神经网络加速比较高的实现方式
- 通过泛函数逼近定理精确地模拟非线性函数



算子级加速：非线性算子加速

• 非线性函数转换：数学表达式近似

GELU

- ❑ 误差函数近似
- ❑ 根据输入范围近似

$$\begin{aligned} \text{erf}(x) &= \frac{2}{\sqrt{\pi}} \int_0^x \exp(-t^2) dt \\ &= \text{sgn}(x)[a(\text{clip}(|x|, \max = -b) + b)^2 + 1] \end{aligned}$$

LayerNorm

- ❑ 迭代实现开根号效果
- ❑ 重复 $x_{i+1} = \left\lfloor \left(x_i + \frac{n}{x_i}\right)/2 \right\rfloor$

Softmax

- ❑ 和最大值进行归一化

$$\begin{aligned} \text{Softmax}(x) &= \frac{\exp(x_i)}{\sum_{j=1}^k \exp(x_j)} = \frac{\exp(x_i - x_{\max})}{\sum_{j=1}^k \exp(x_j - x_{\max})} \\ &= \exp[x_i - x_{\max} - \ln(\sum_{j=1}^k \exp(x_j - x_{\max}))], \end{aligned}$$

缩小输入范围

- ❑ 指数函数的多项式近似

$$\exp(x) = 0.3585(x + 1.353)^2 + 0.344$$

避免除法器

- ❑ 指数函数/对数函数拆分2的基数

$$e^y = 2^{u+v} = 2^u \times 2^v = \begin{cases} 2^v \ll u & u > 0 \\ 2^v \gg u & u \leq 0 \end{cases}, u = \lfloor y \times \log_2 e \rfloor, v = y \times \log_2 e - u$$

小数运算+移位

$$\ln F = \ln 2 \times \log_2 F = \ln 2(w + \log_2 k) = \ln 2(k - 1 + w), F = 2^w + k$$

目录

- 研究背景
- 模型级加速
 - 模型压缩
 - 稀疏性处理
- 模块级加速
 - 注意力模块加速
 - 前馈神经网络加速
- 算子级加速
 - 线性算子加速
 - 非线性算子加速
- 总结及展望

北京大学-智能硬件体系结构

2024年秋季学期

主讲：陶耀宇、李萌

总结及展望



目录

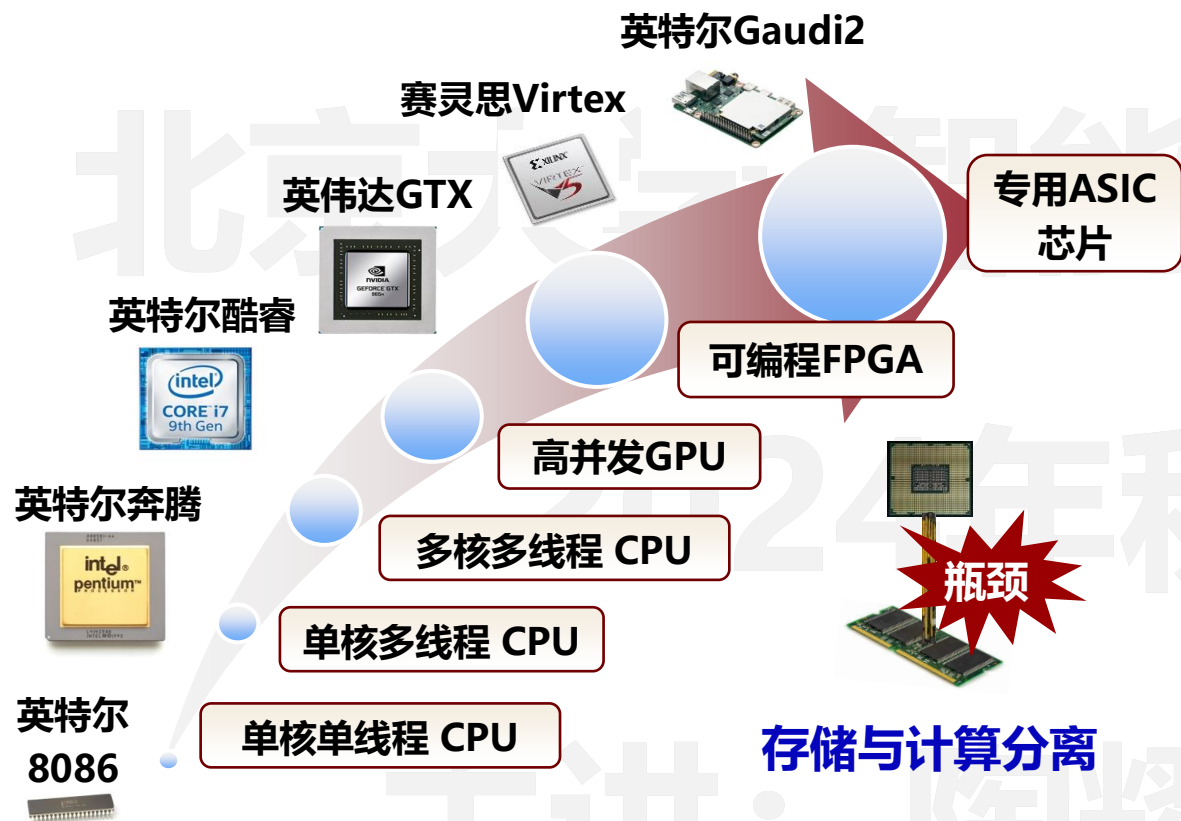
CONTENTS



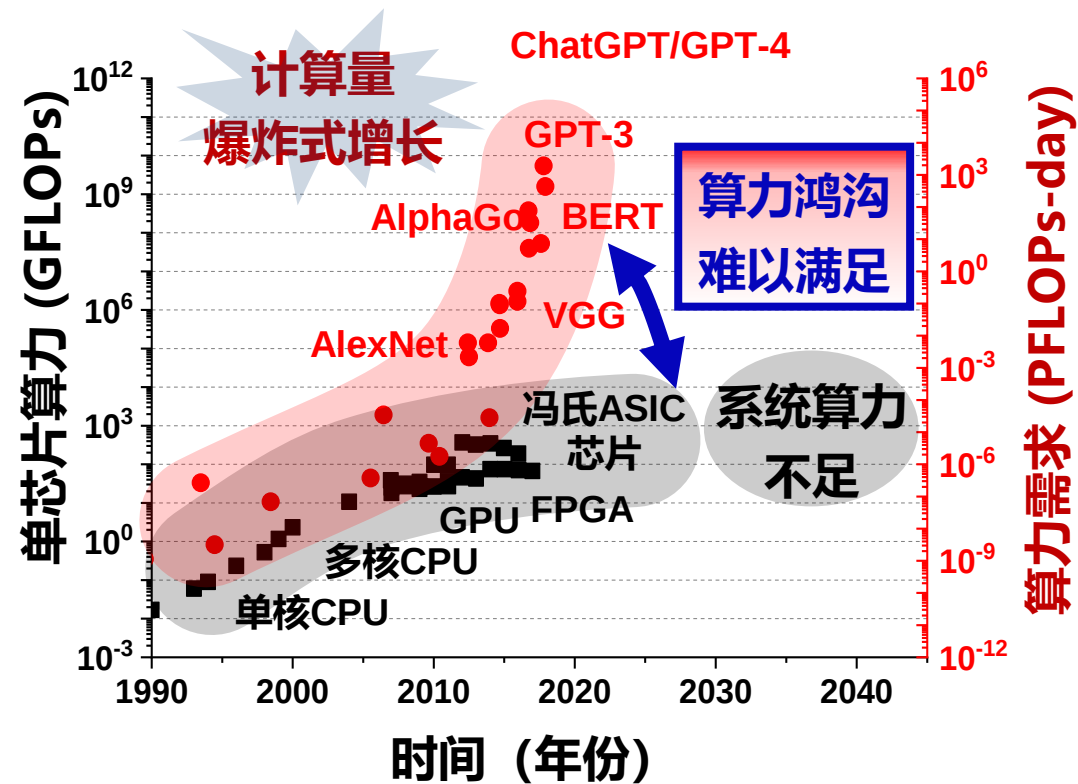
- 01. 典型AI芯片架构**
- 02. AI芯片软硬件协同设计**
- 03. AI大模型芯片设计**
- 04. 存算一体AI芯片架构**

当前半导体芯片产业发展的关键瓶颈

- 以冯诺依曼架构芯片为基石的传统算力系统难以满足爆炸式的算力增长需求



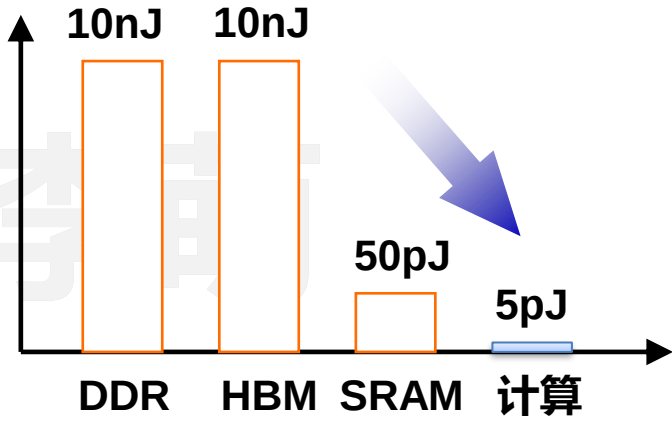
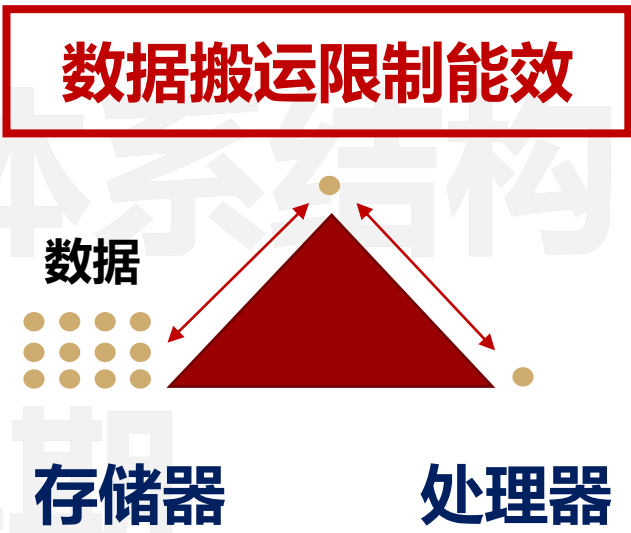
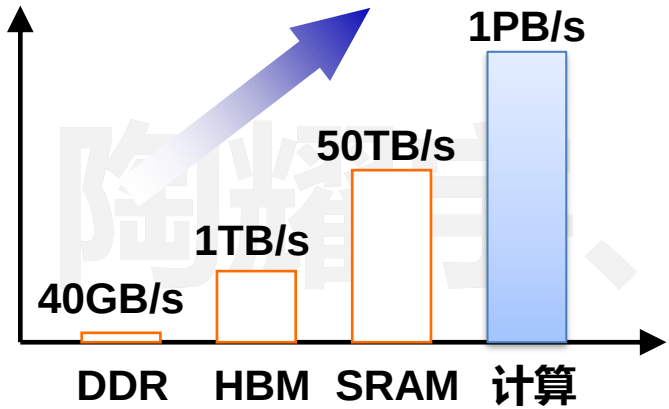
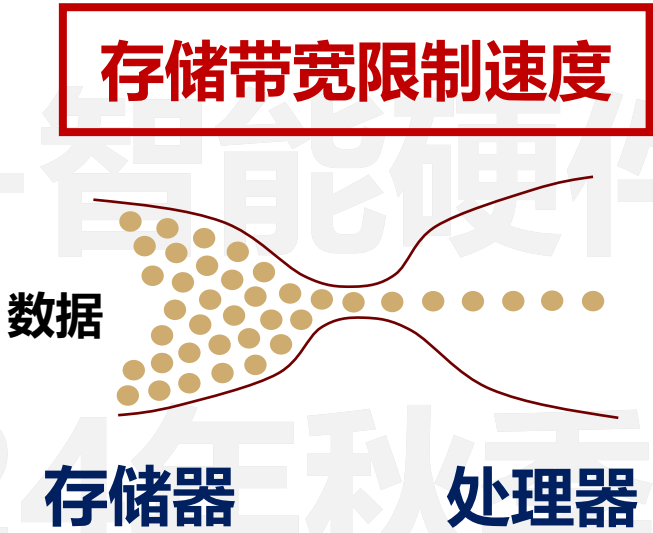
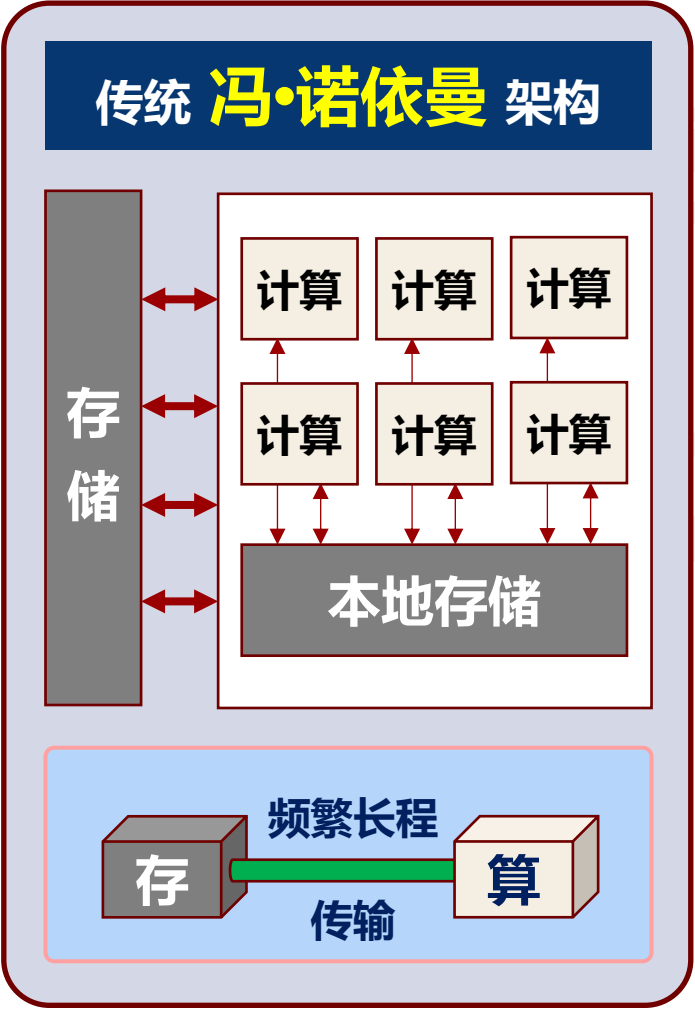
摩尔定律时代的冯诺依曼架构芯片演进图



新兴计算任务所需的计算量呈爆炸式增长

当前半导体芯片产业发展的关键瓶颈

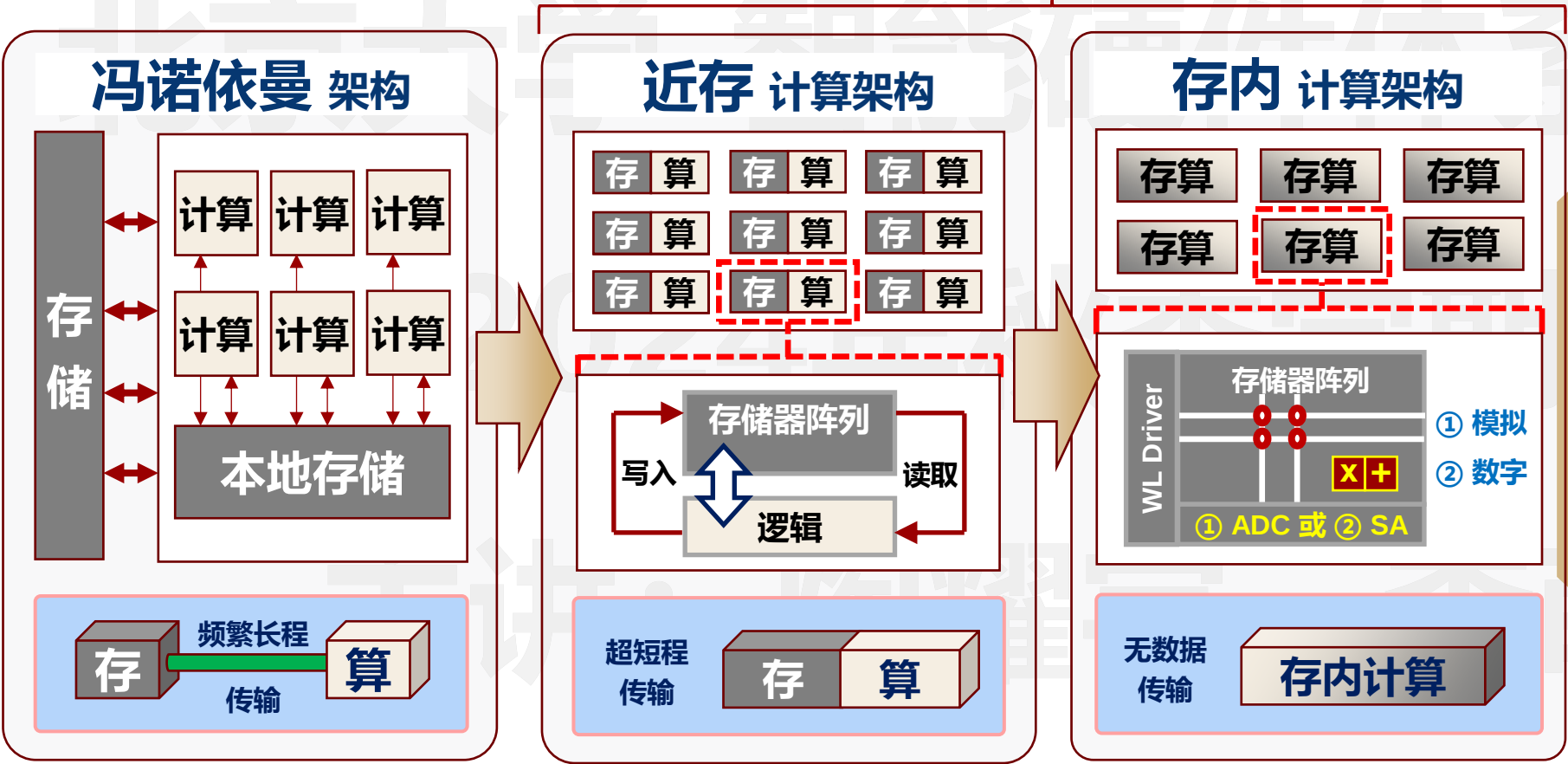
- 传统冯式架构的计算芯片难以突破大算力和高能效的性能瓶颈



代表性半导体芯片前沿技术：存算一体

- 存算一体技术成为后摩尔时代打破算力瓶颈的重要路径

算力提升、能效提升 → 存算一体技术

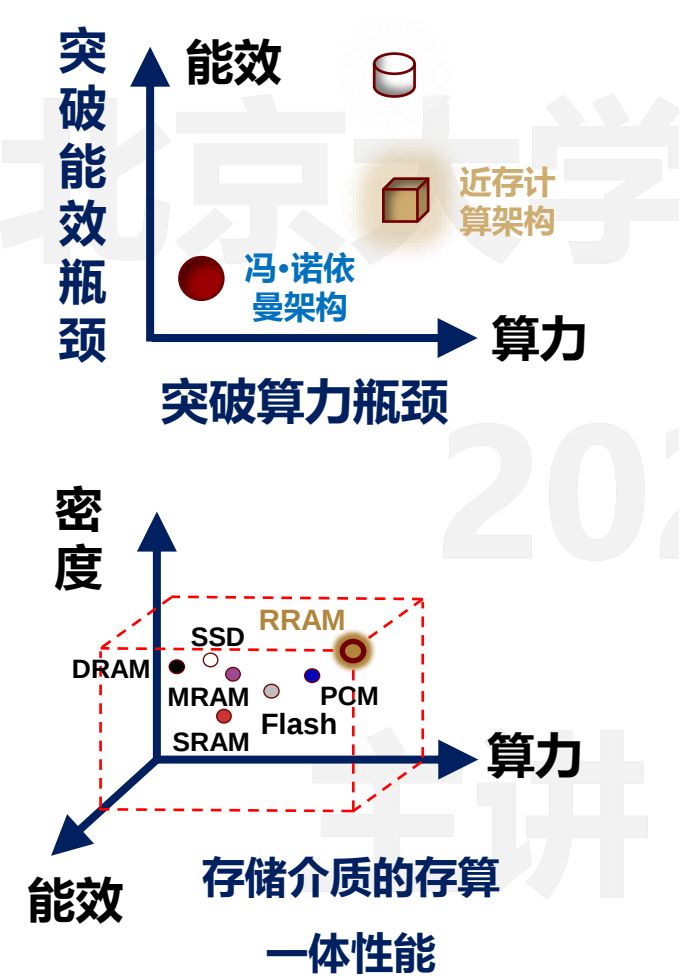


天然优势

- 大算力
- 低功耗
- 低延时

实现存算一体的多种存储介质

- ReRAM是最具发展潜力的新兴存储器之一，非常适合存算一体



存储器	主要优势	缺点与不足	非易失	适合场景
DRAM	工艺成熟 密度高	速度低/刷新/ 只能做近存	否	冯氏架构 过渡
SRAM	工艺成熟 IP化	能效低/密度低	否	端侧/边缘 中小算力
SSD	容量大 成本低	速度低/近存	是	云端大容量
MRAM	能效/速度/ 密度高	大规模集成难	是	端侧/边缘 中小算力
Flash	密度高/ 成本低	PVT变化敏感	是	端侧/边缘 低成本
ReRAM	算力/能效/ 密度高	良率爬坡	是	端侧/边缘 大算力

存算一体在国内外产业界和学术界的布局

- 存算一体技术受到国内外产业界和学术界的高度重视，已有诸多产业化布局



实现存算一体的多种技术路径

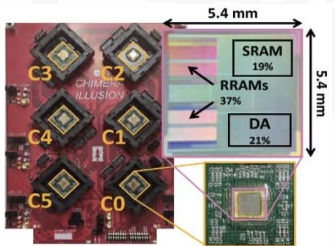
- 存算一体技术受到国内外产业界和学术界的高度重视，已有诸多产业化布局

器件

基于嵌入式存储的近存计算

存内没有计算

多种嵌入式存储器
(DRAM、SRAM
等)

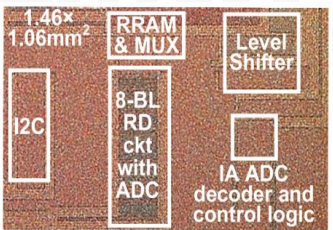


M. Giordano et al., VLSI 2021

基于新兴非易失存储器的存内计算

存内含计算

需新兴器件与电路
架构的协同创新



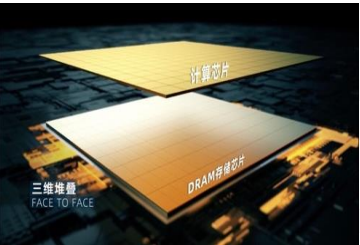
J.H. Yoon et al., JSSC 2022

计算架构

基于先进封装的近存计算

存内没有计算

多种先进封装
(Chiplet、三维
封装等)

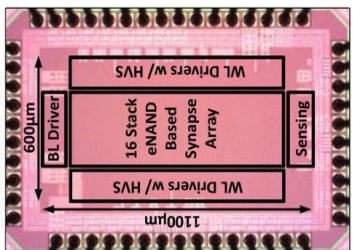


D. Niu et al., ISSCC 2022

基于存储级存储器的存内计算

存内含计算

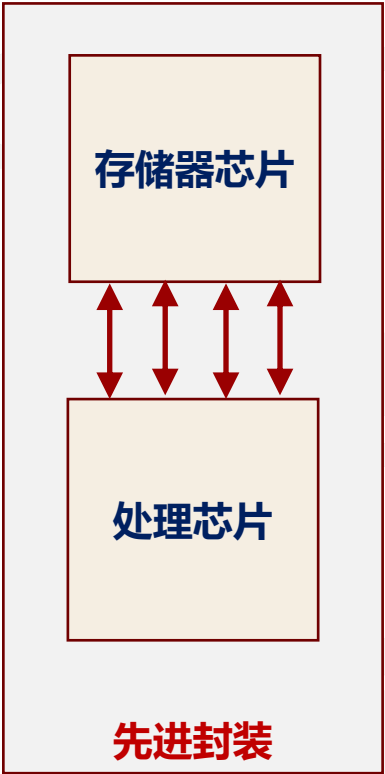
将存内计算应用于
成熟的存储级存储
器 (如Flash)



M. Kim et al., JSSCC 2022

路径一：先进封装实现近存计算

- 将存储器芯片与逻辑处理芯片通过多种先进封装方式拉近距离



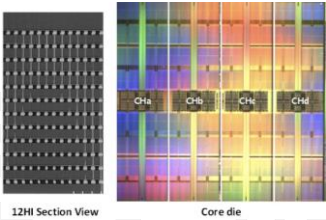
优势

与现有系统兼容性好
设计难度低

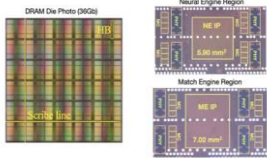
典型案例

AMD Zen CNM, SK-Hynix
HBM-PIM, Alibaba CNM

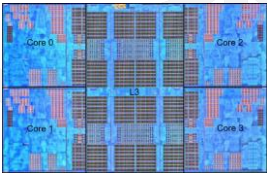
性能	先进封装 近存计算
算力	高
能效	低
设计难度	低
可重构性	高
成本	低



SK-海力士
HBM-PIM



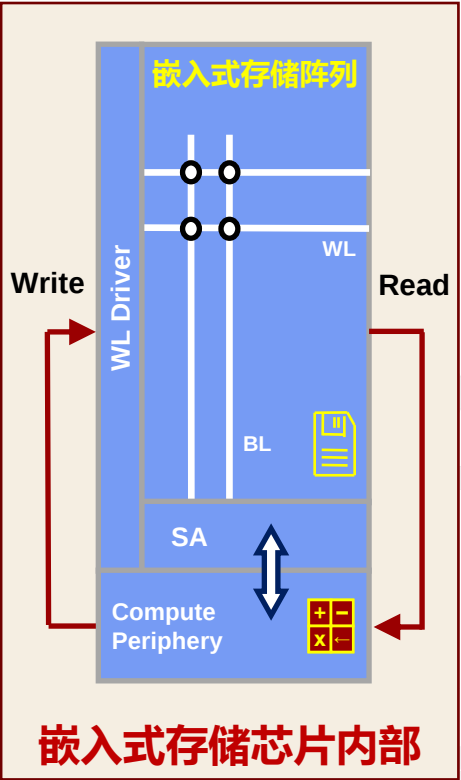
阿里巴巴达摩院3D
封装近存计算芯片



AMD Zen近存
计算样片

路径二：嵌入式存储实现近存计算

- 在嵌入式存储器芯片（DRAM/SRAM等）内部集成近存逻辑电路



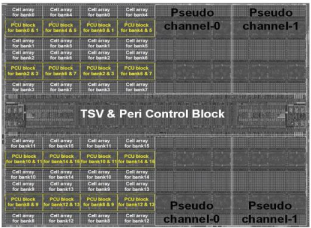
优势

数据搬运距离较先进封装
更加缩短

性能	嵌入式存储近存
算力	中等
能效	中等
设计难度	中等
可重构性	中等
成本	中等

典型案例

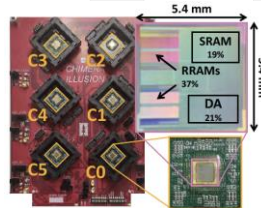
三星 HBM CNM,
GraphCore SRAM CNM,
Stanford RRAM CNM



三星HBM CNM



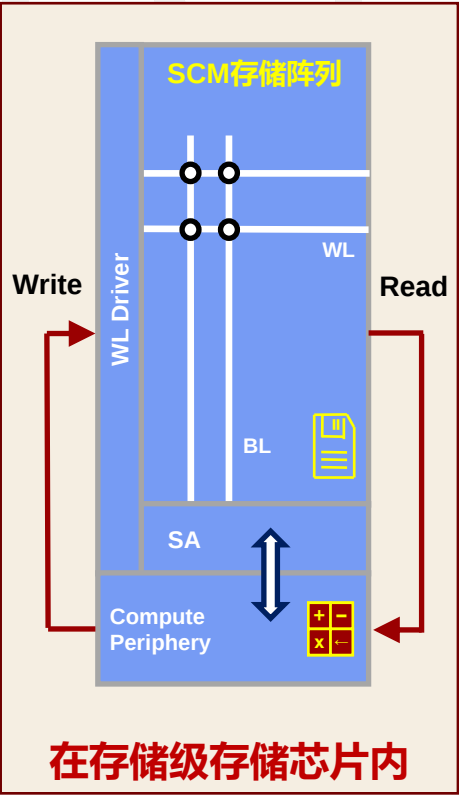
GraphCore SRAM CNM



Stanford RRAM CNM

路径二：存储级存储实现存内计算

- 利用成熟的存储级存储器（Flash等）进行存内计算



优势

工艺成熟、成本低

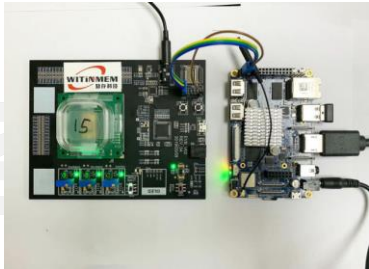
典型案例

Mythic Flash CNM,
WitinMem Nor Flash CIM

性能	存储级存储器 存内计算
算力	中等
能效	低
设计难度	中等
可重构性	中等
成本	低



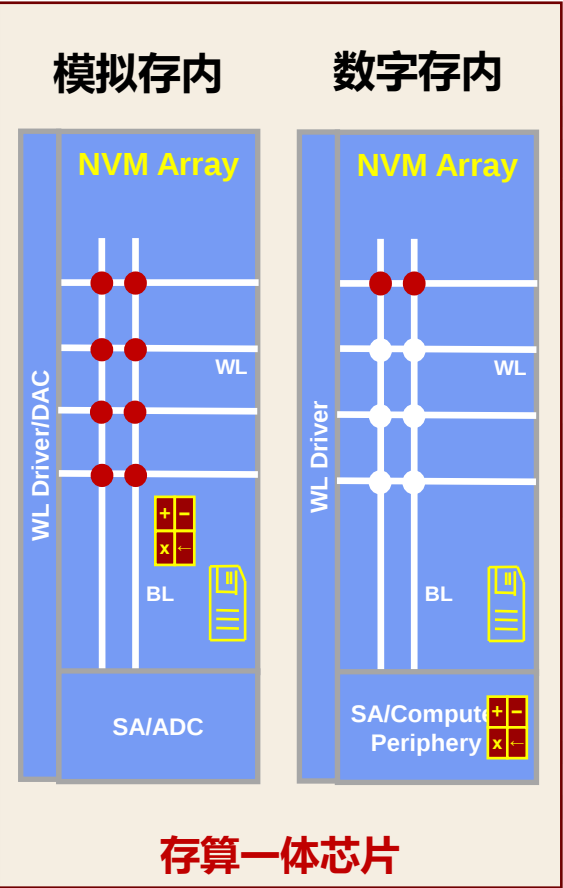
Mythic Flash CIM Chip



知存科技 Flash CIM

路径二：新兴非易失性器件实现存内计算

- 使用新兴的非易失性存储器（ReRAM等）进行存内计算



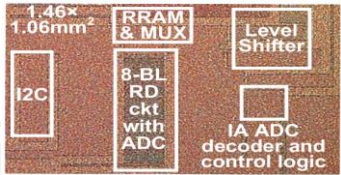
优势

大算力、高能效、高密度

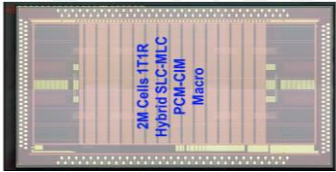
典型案例

台积电ReRAM芯片、台积电PCM芯片、三星MRAM芯片

性能	新兴非易失性存储器存内计算
算力	高
能效	高
设计难度	高
可重构性	中等
成本	高



TSMC RRAM CIM



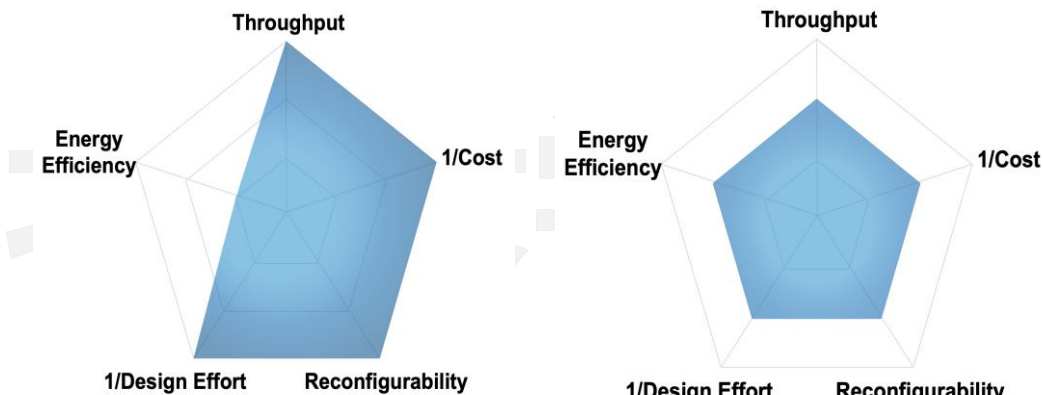
TSMC PCM CIM



三星 MRAM CIM

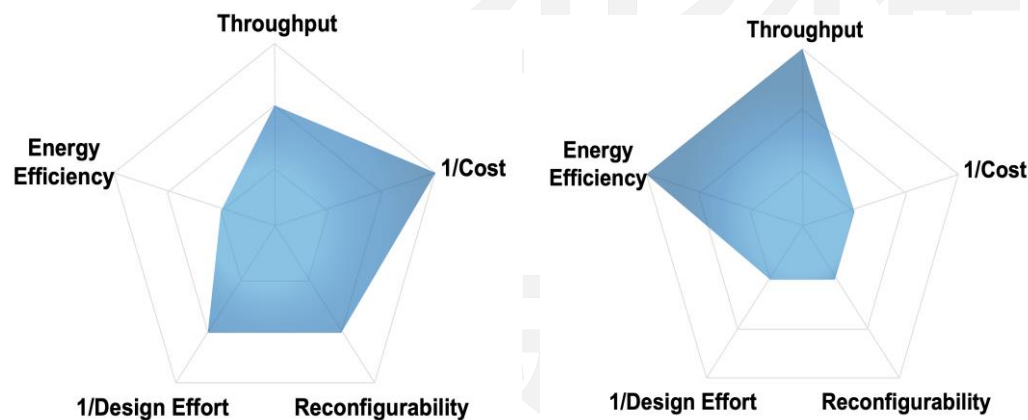
存算一体多种技术路径的比较

- 存算一体技术针对不同应用场景可以采用不同的技术路径



先进封装近存 CNM-3D

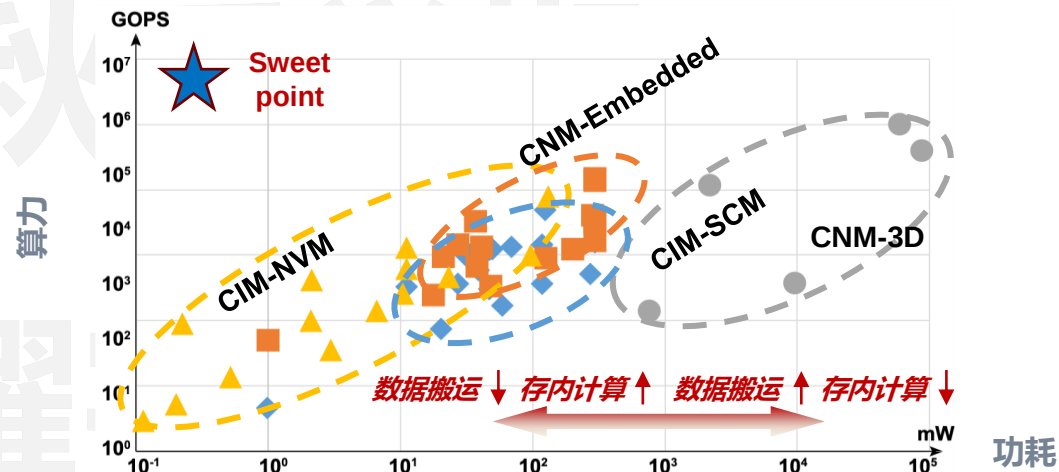
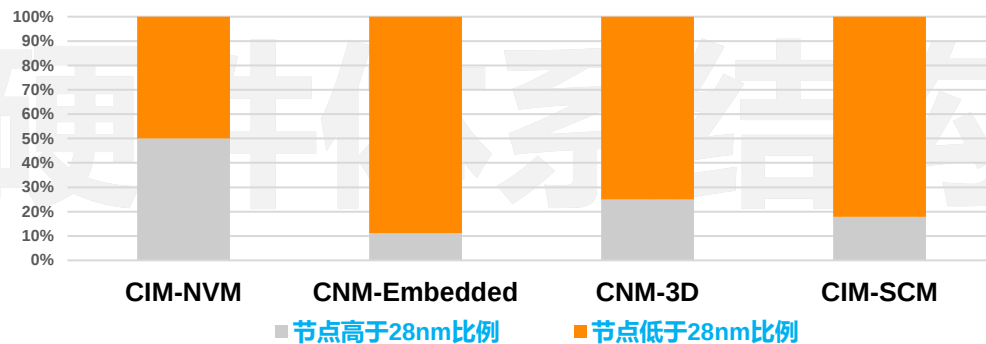
嵌入式近存 CNM-Emb.



SCM近存/存内 CIM-SCM

新兴NVMs存内 CIM-NVM

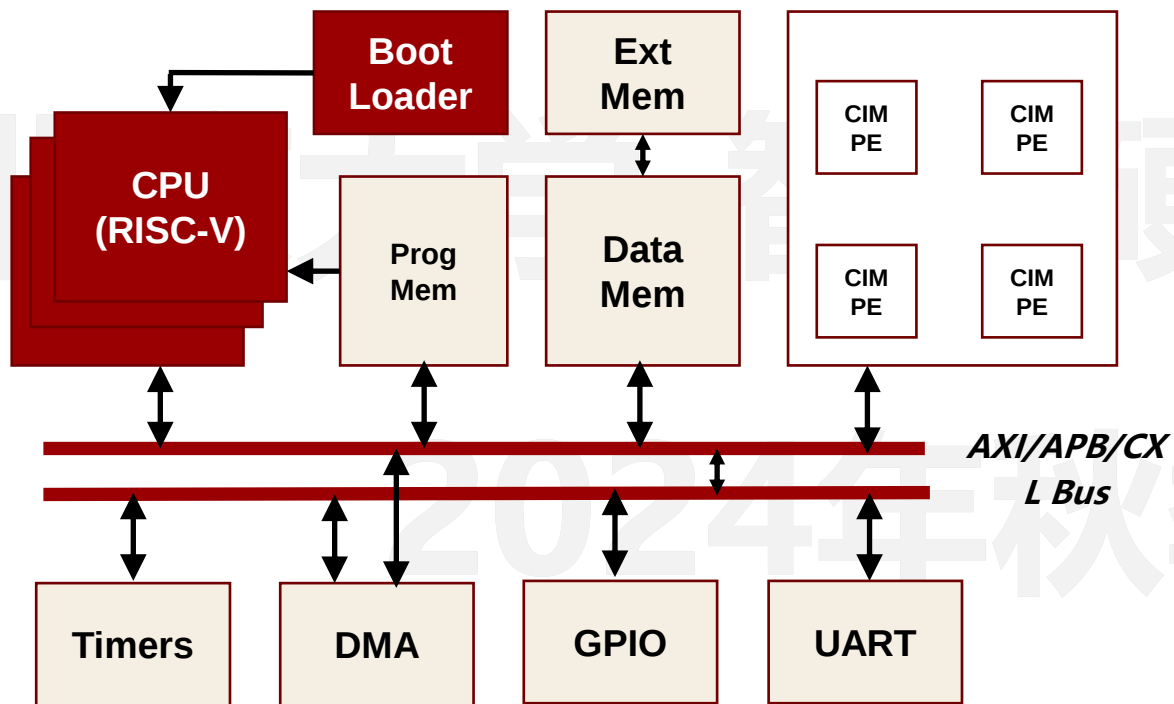
新兴NVMs存内计算 (CIM-NVM) 向先进节点演进中



当前存算一体芯片算力、功耗图谱

存算一体的异构集成

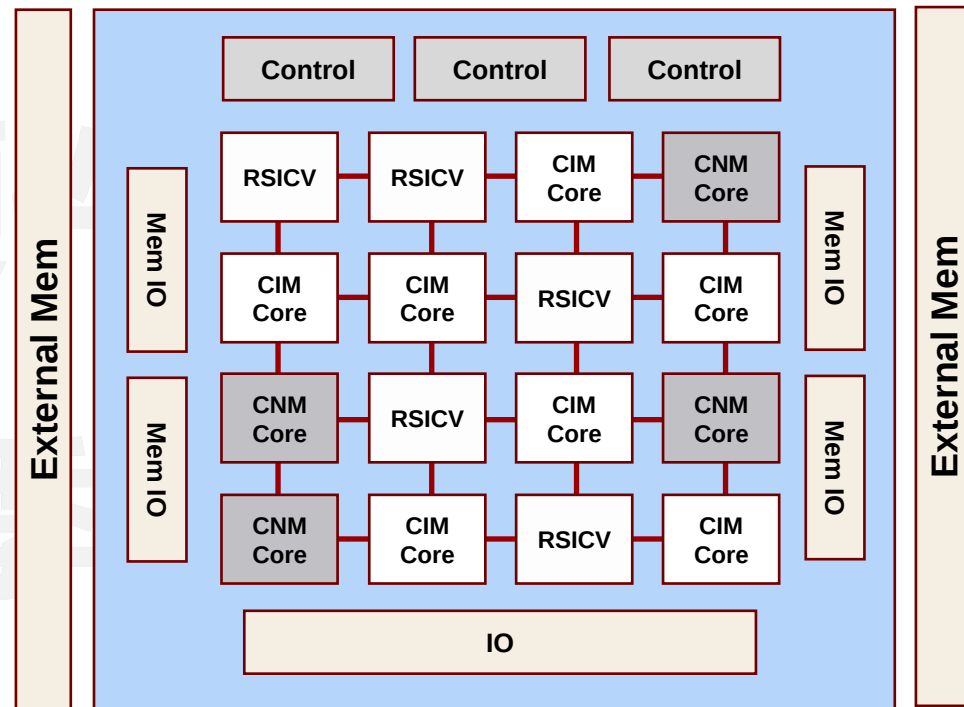
- 融合存算一体技术的多阵列异构集成芯片



集中式异构集成架构

设计难度相对较低

但基于集中式总线的架构难以最大限度发挥存算一体的优势



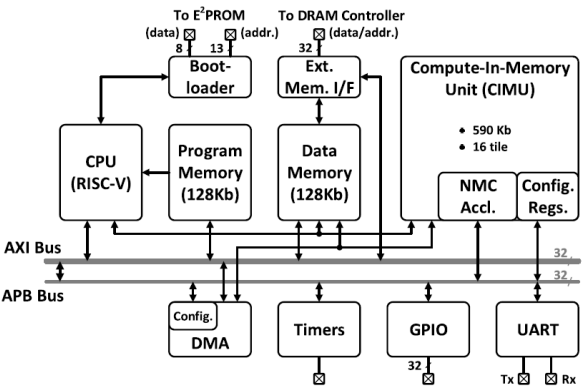
分布式异构集成架构

可扩展性、可重构性好

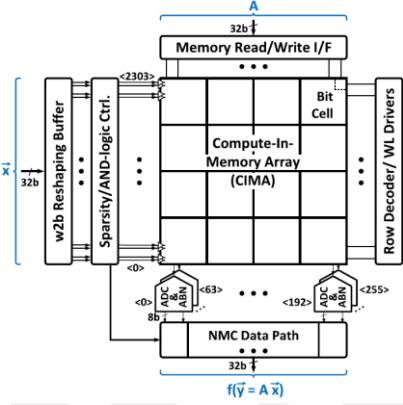
但片上网络设计及通信难度高

典型存算一体异构集成芯片 – 普林斯顿总线式集成芯片

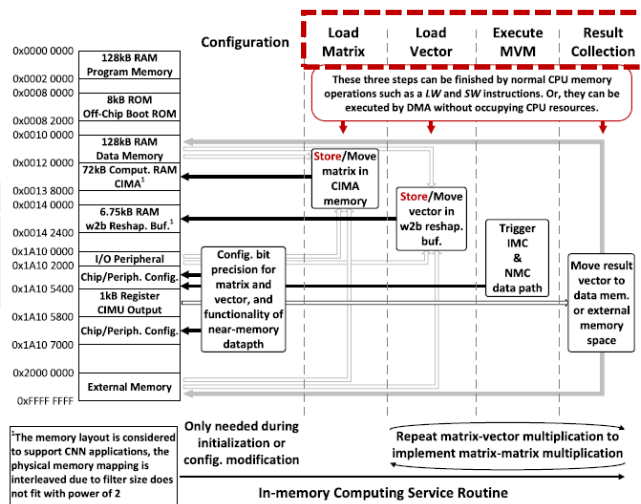
融合存算一体技术的多阵列异构集成芯片



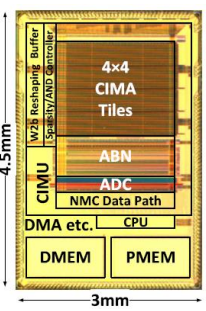
顶层架构



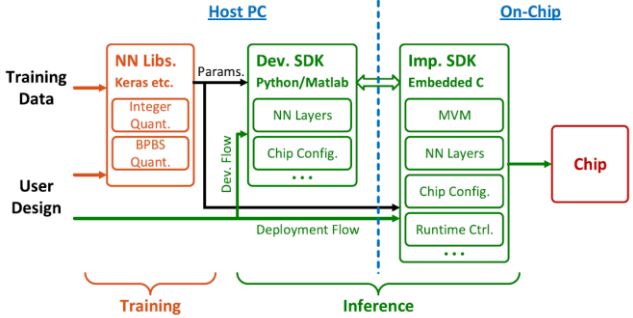
存内计算阵列架构



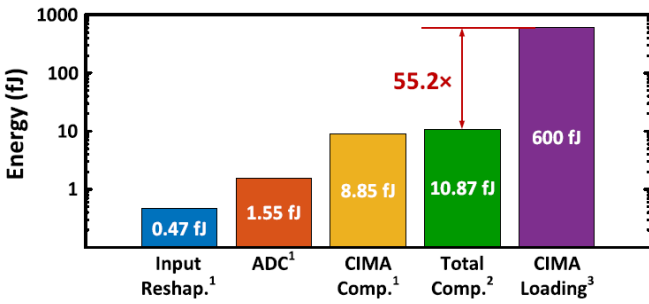
存内计算实现矩阵乘法时序图



芯片版图



存内和非存内计算的映射方式

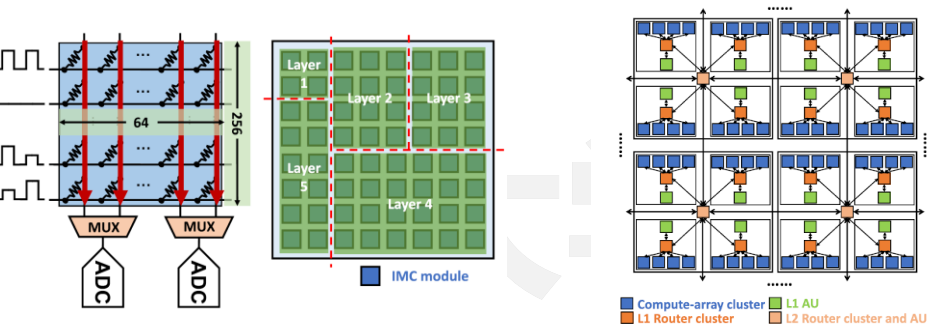


能耗分布分析

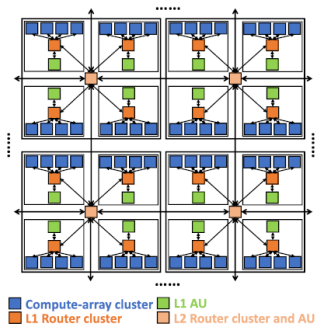
- 4x4 8T SRAM 存内计
- 阵列单阵列576x64, 共 590Kb
- 集成RISCV CPU
- 128Kb ProgMem
- 128Kb DataMem
- 模拟存内 8b ADC 精度

典型存算一体异构集成芯片 – 密歇根大学TAICHI分布式集成芯片

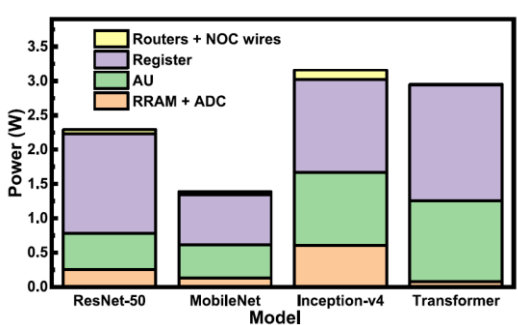
融合存算一体技术的多阵列异构集成芯片



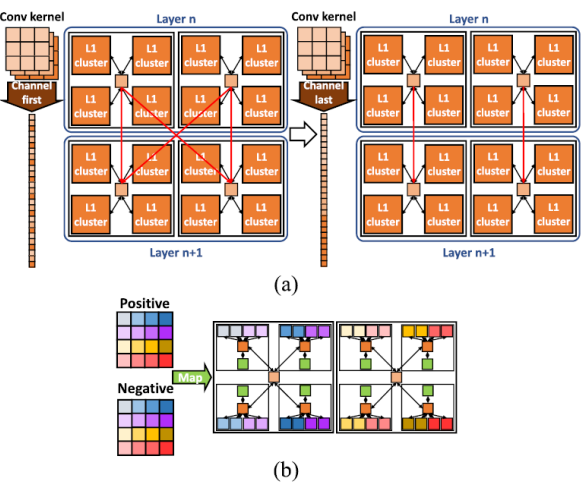
ReRAM模拟存内计算阵列



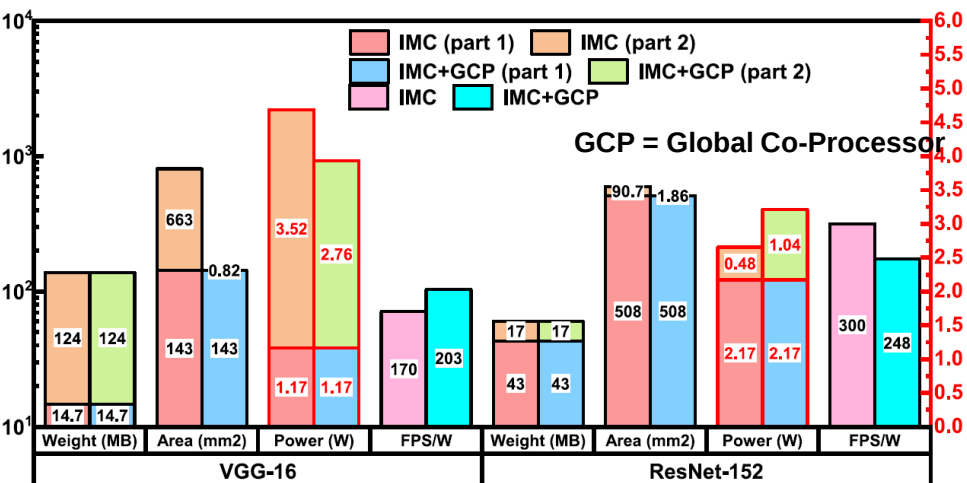
层次式片上网络



能耗分布分析



4x4个cluster、每个cluster包含4x4个存内计算阵列

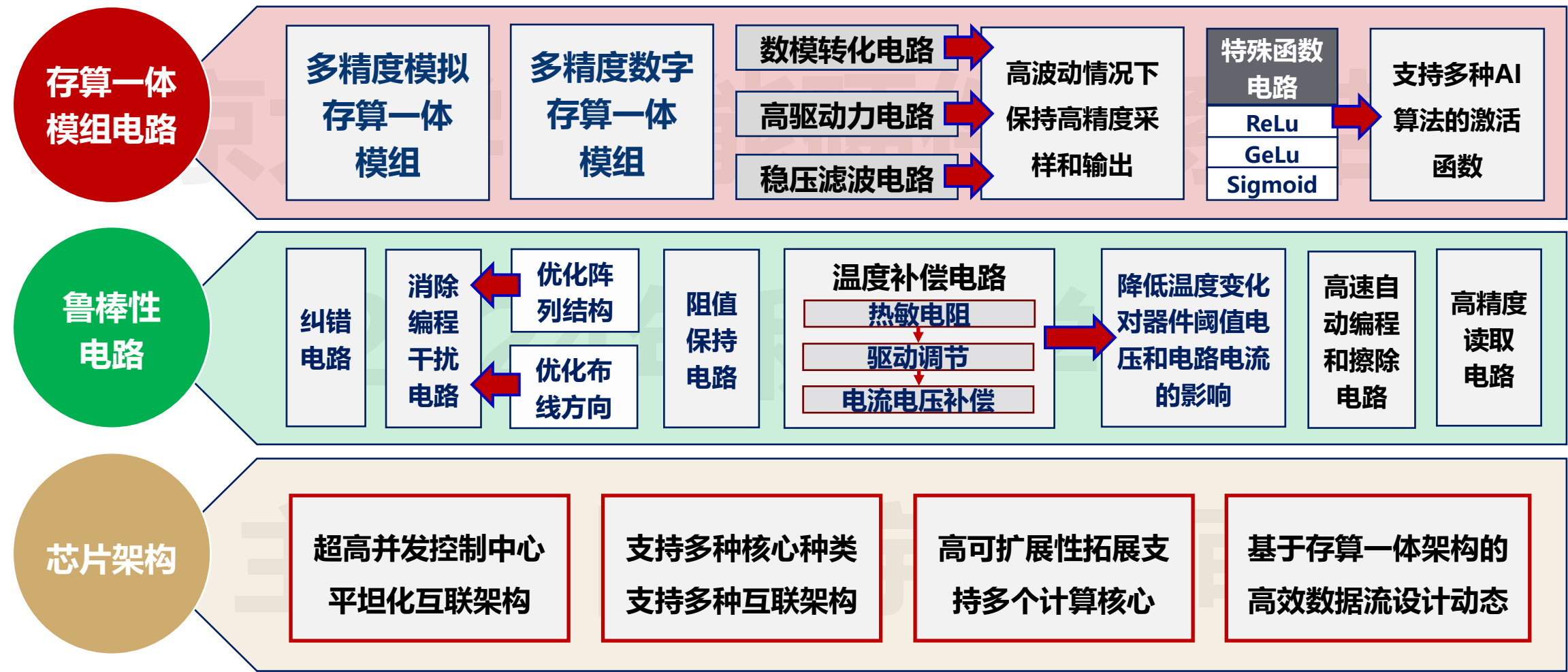


分布式异构架构运行VGG-16和ResNet-152的性能分析

- 16x16层级式网格片上网络
- 集成4种不同尺寸的ReRAM
- 模拟存内计算阵列
- 模拟存内 8b ADC 精度

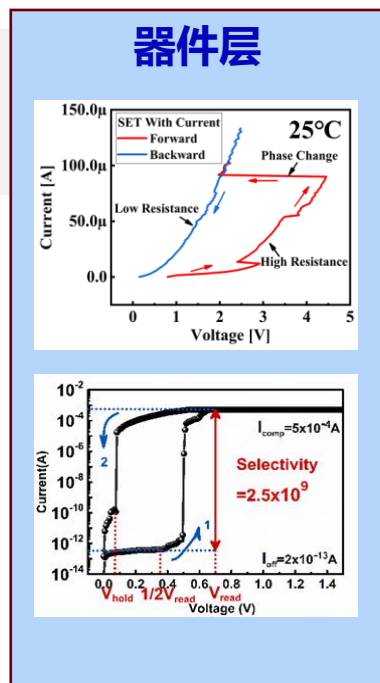
存算一体异构集成的电路与架构挑战

- 优化电路与芯片架构，保障能效优势和迭代演进能力

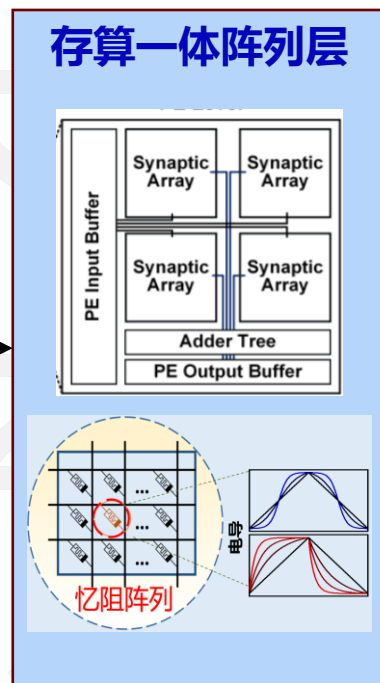


存算一体异构集成的系统级多层次挑战

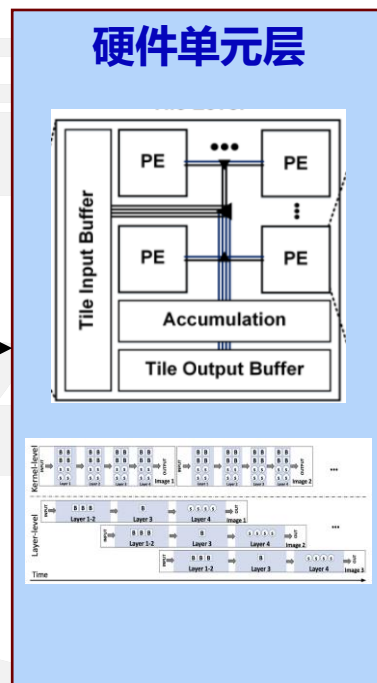
- 需要考虑多层次的协同设计，设计空间非常大



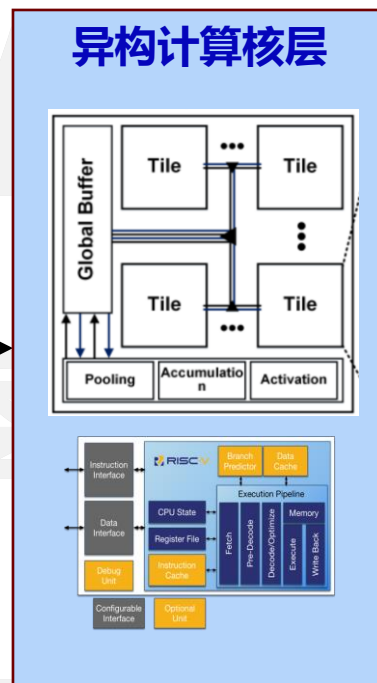
Non-Ideality, etc.



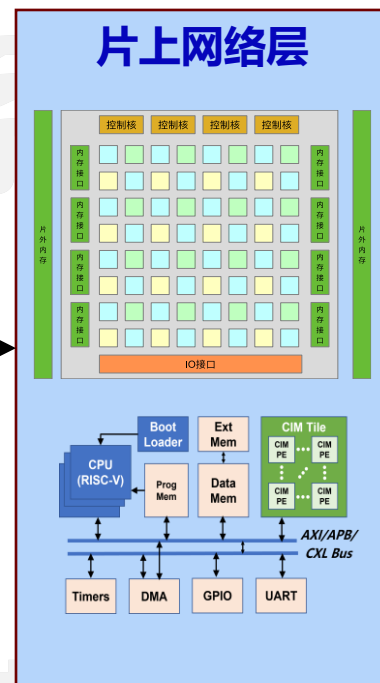
BER, ADC, etc.



PE Pipeline, etc.



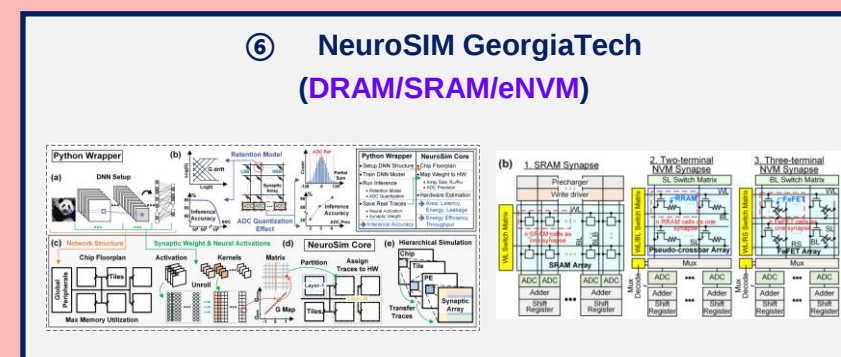
Core Configs, etc.



NoC/Bus Traffic, etc.

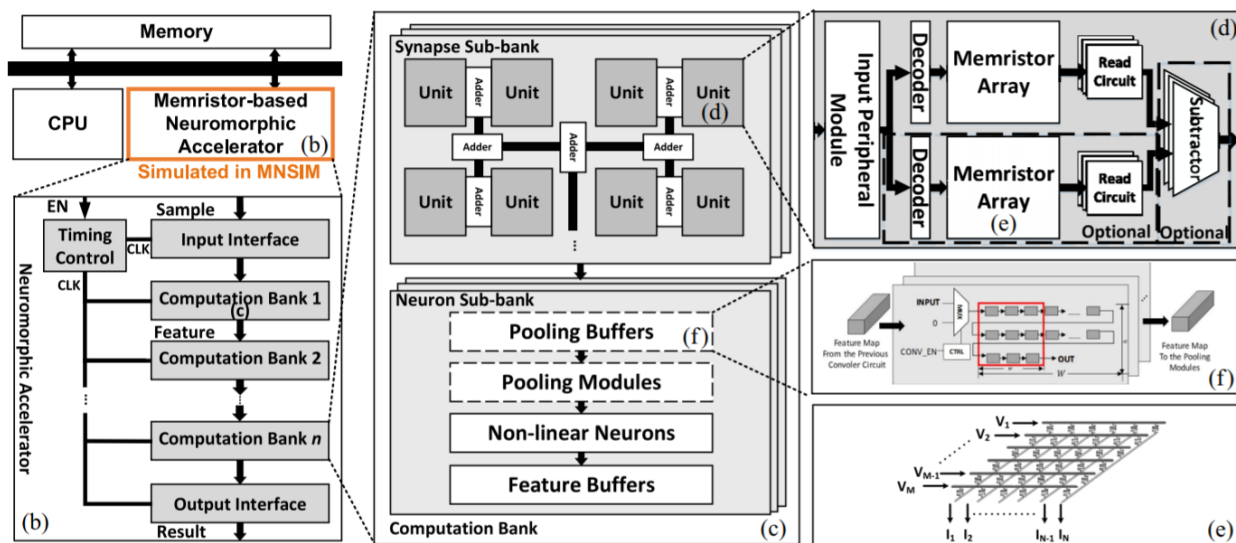
- ## 目标应用

AI计算 (MAC/Conv)

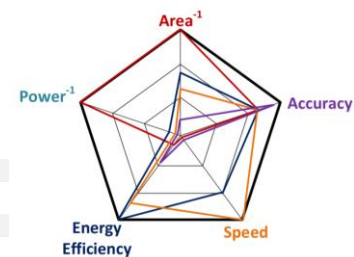


典型存算一体架构仿真器 – 清华MNSIM

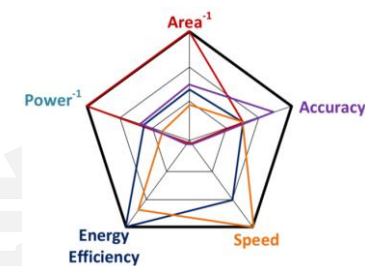
• 面向AI计算的存算一体架构仿真器



Hierarchical structure: bank-level, tile level and PE-level



Deep CNN simulation demos

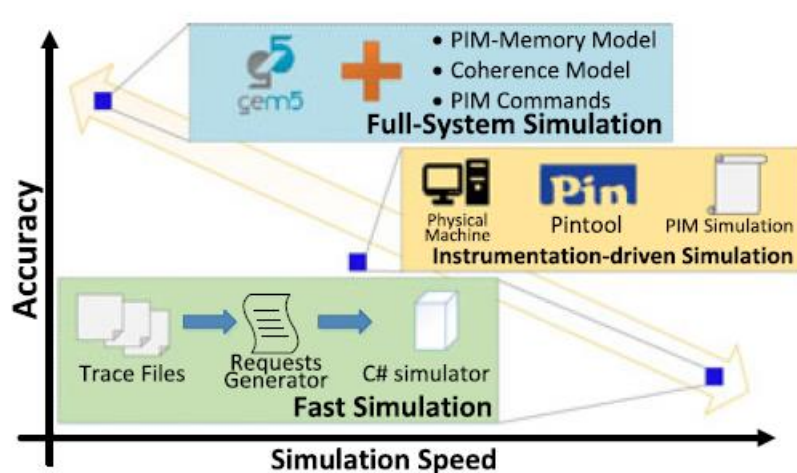


缺乏器件级的仿真

- 层级式的RRAM AI计算仿真器 – 覆盖算法、电路、架构、系统多个层级
- 灵活的接口配置 – 在各个层级接口均可以进行灵活配置，搜索空间很大
- 搜索速度快 – 仿真器均采用准确的行为级模型，可以并行仿真加快仿真速度

典型存算一体架构仿真器 – 南加州大学PIMSIM

• 面向通用计算的存算一体架构仿真器



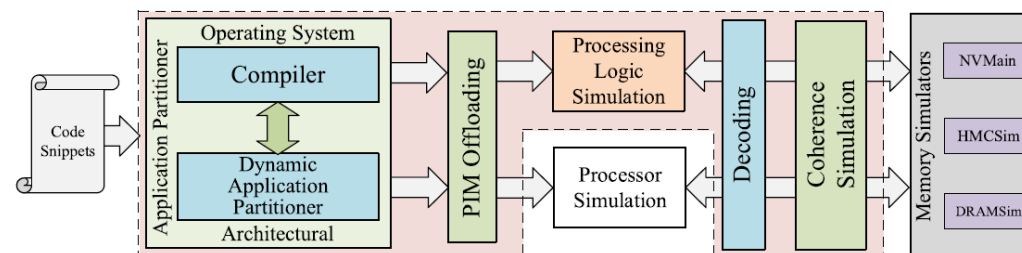
Three Simulation Modes

- ① Full-System
- ② Instrumentation-driven
- ③ Fast

```
1: #PIM_START
2: #ID=0 //define the execution logic ID
3: #INPUT = a
4: #INPUT = b
5: #OUTPUT = c
6: c = a + b;
7: #PIM_END
```

```
1: #PIM_CODE_ADD
2: #ID=0 //define the execution logic ID
3: #DEFINE_INPUT Input1
4: #DEFINE_INPUT Input2
5: #DEFINE_OUTPUT Output1
6: Output1 = Input1 + Input2;
7: #PIM_END
8: PIM_ADD(c, a, b);
```

Need to
manually
label PIM
tasks



PIMSIM Architecture Simulation Framework

- 提供三种仿真模式 – Tradeoff between simulation speed and accuracy
- PIM任务标识 – 可执行的PIM任务在汇编编译的阶段进行代码标识
- 兼容各类DRAM后端仿真器 – 例如DRAMSim等

- 课程评估：扫右侧二维码或登录网上评估系统
(kcpkg.pku.edu.cn)
 - 希望同学们能给予好评，感谢！
- 希望大家后续科研学习一切顺利，也欢迎在课程结束后，同学们仍然能找我讨论有趣的问题
- 办公室：资源西楼2208B
- 邮箱：taoyaoyutyy@pku.edu.cn
- 电话：13095698730（微信同）

