

**高效注意力**

**清华大学计算机系 陈键飞**

**《大模型计算》课程团队**

# 计算量/访存量

❖ 考虑8B模型: seqlen=8192, head\_dim=128

$$S = q * k.t() / \text{sqrt}(d)$$

计算量 $8192*8192*128*2/2 = 8.59\text{Gflops}$  (causal)

访存量 $8192*128*2*2\text{bytes} + 8192*8192*4\text{bytes} = 0.254\text{GB}$

$$P = \text{softmax}(S)$$

计算量可忽略不计

访存量 $8192*8192*2*4\text{bytes} = 0.5\text{GB}$  (一读一写)

$$O = P * v$$

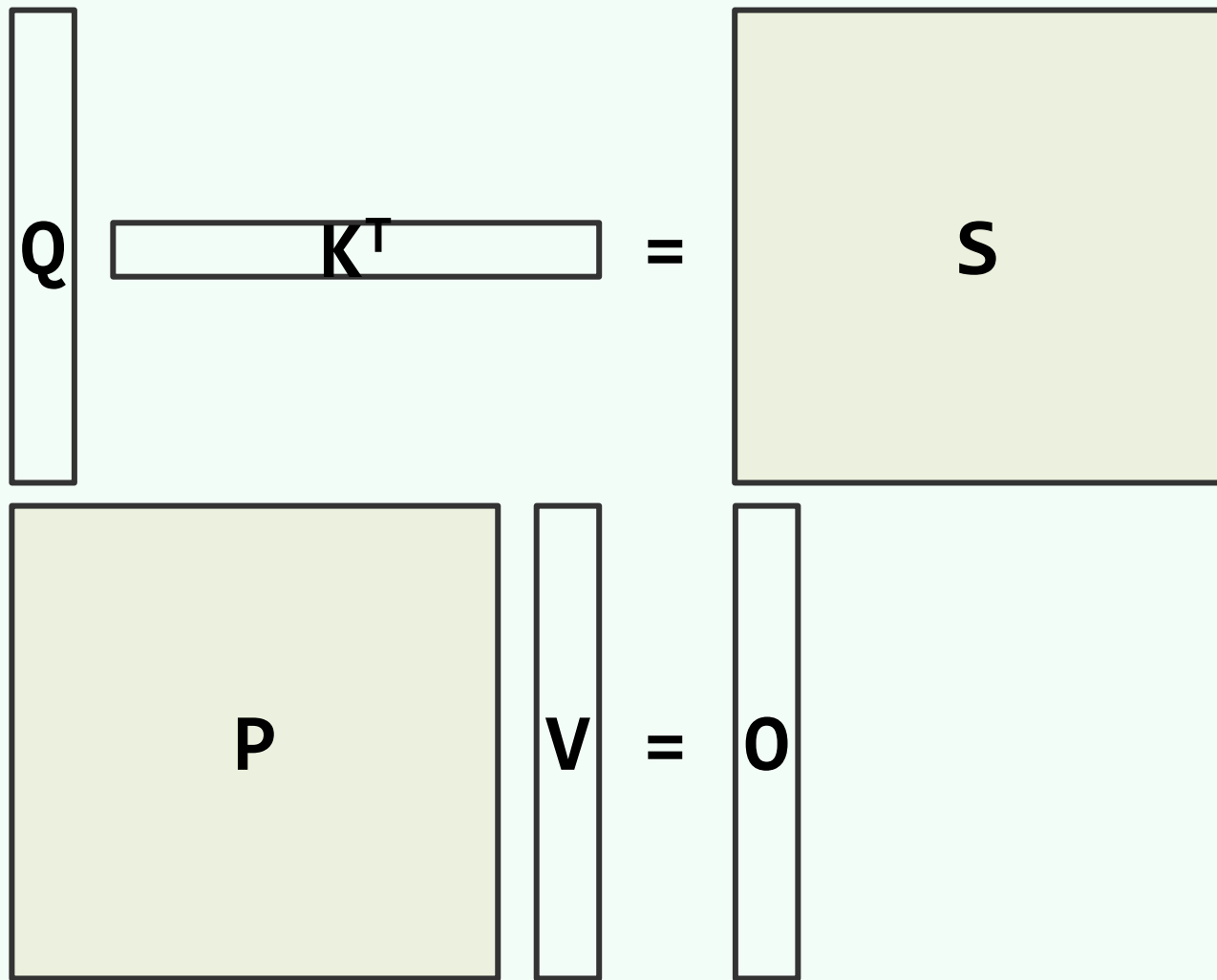
计算访存量与计算S相等

计算量=17.2G, 访存量=1.008GB, **计算访存比=17 flops/Byte**

在RTX 4090上, 如果打满1008GB/s带宽 (FP16), 则输出算力=17.2Tflops

理论峰值165Tflops

# 问题



算子的输入输出 (QKVO) 都很小  
但中间结果 (SP) 很大

# 思路

- ❖ QKV一起往GPU里喂
- ❖ SP不出芯片，直接吐出O

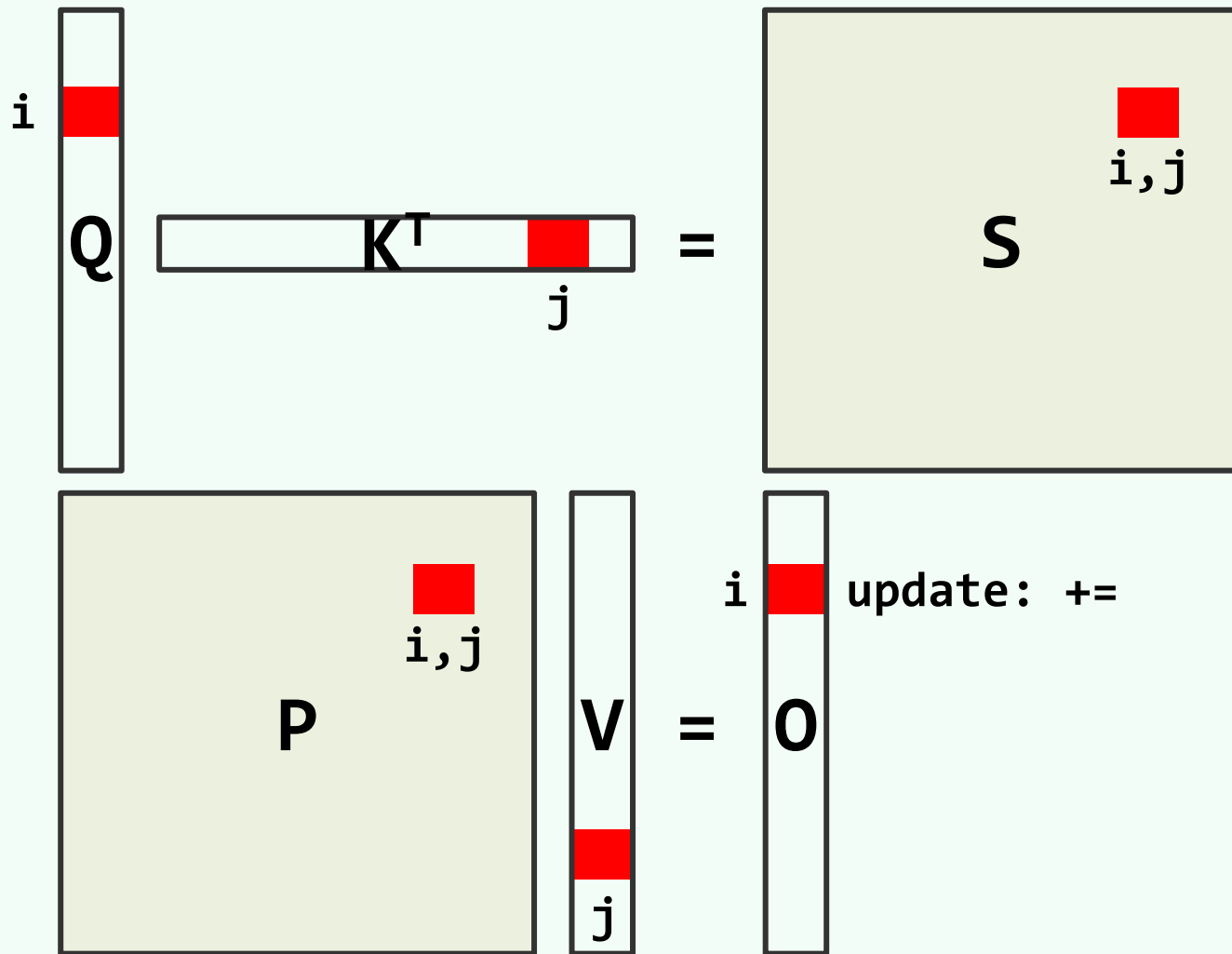
$$O_i += P_{ij} V_j$$

$$P_{ij} = ? (S_{ij})$$

$$S_{ij} = Q_i K_j^T$$

唯一的问题：softmax

$P_{ij}$  不仅取决于  $S_{ij}$ ，还取决于其他tile



## Online softmax

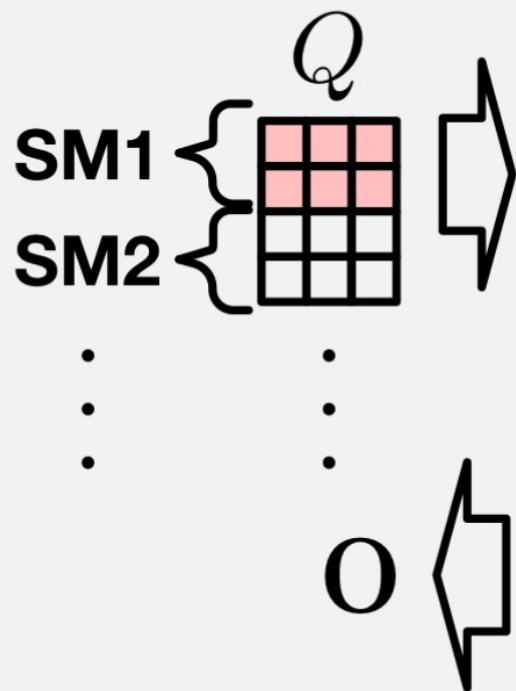
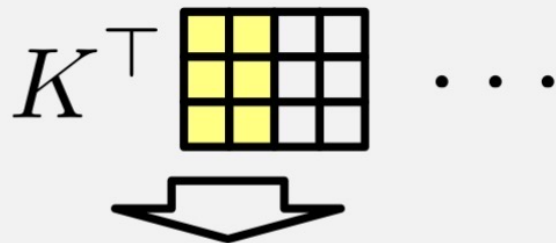
- ❖ 要算  $O = \sum_j P_j V_j = \sum_j \left( \frac{\exp(S_j)}{\sum_k \exp(S_k)} \right) V_j$
- ❖ 问题:  $\sum_k \exp(S_k)$  事先不知道
- ❖ 解法: 分成两个求和  $O = \frac{(\sum_j \exp(S_j) V_j)}{\sum_k \exp(S_k)}$
- ❖ 循环元素j: 维护两个求和  $O += \exp(S_j) V_j$  和  $\ell += \exp(S_j)$
- ❖ 最后输出  $O / \ell$
- ❖ 期间  $O$  和  $\ell$  可以同时缩放避免softmax带来的数值问题
- ❖ 额外维护  $m = \max\{m, S_j\}$
- ❖ 将两个求和同时缩放到  $m$ :  $O = O / \exp(m^{new} - m^{old})$ ,  $\ell = \ell / \exp(m^{new} - m^{old})$

# FlashAttention

Loop for  $K$  blocks

Global Memory

SM

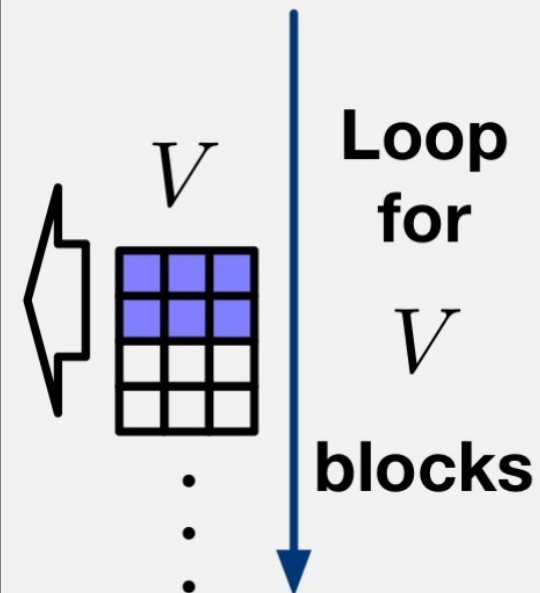


SM1

$$S = (\text{red grid} \cdot \text{yellow grid})$$

$$\tilde{P} = \text{OnlineSoftmax}(S)$$

$$O = (\tilde{P} \cdot \text{blue grid})$$



---

**Algorithm 1** FLASHATTENTION-2 forward pass

---

**Require:** Matrices  $\mathbf{Q}, \mathbf{K}, \mathbf{V} \in \mathbb{R}^{N \times d}$  in HBM, block sizes  $B_c, B_r$ .

- 1: Divide  $\mathbf{Q}$  into  $T_r = \left\lceil \frac{N}{B_r} \right\rceil$  blocks  $\mathbf{Q}_1, \dots, \mathbf{Q}_{T_r}$  of size  $B_r \times d$  each, and divide  $\mathbf{K}, \mathbf{V}$  into  $T_c = \left\lceil \frac{N}{B_c} \right\rceil$  blocks  $\mathbf{K}_1, \dots, \mathbf{K}_{T_c}$  and  $\mathbf{V}_1, \dots, \mathbf{V}_{T_c}$ , of size  $B_c \times d$  each.
  - 2: Divide the output  $\mathbf{O} \in \mathbb{R}^{N \times d}$  into  $T_r$  blocks  $\mathbf{O}_i, \dots, \mathbf{O}_{T_r}$  of size  $B_r \times d$  each, and divide the logsumexp  $L$  into  $T_r$  blocks  $L_i, \dots, L_{T_r}$  of size  $B_r$  each.
  - 3: **for**  $1 \leq i \leq T_r$  **do**
  - 4:   Load  $\mathbf{Q}_i$  from HBM to on-chip SRAM.
  - 5:   On chip, initialize  $\mathbf{O}_i^{(0)} = (0)_{B_r \times d} \in \mathbb{R}^{B_r \times d}, \ell_i^{(0)} = (0)_{B_r} \in \mathbb{R}^{B_r}, m_i^{(0)} = (-\infty)_{B_r} \in \mathbb{R}^{B_r}$ .
  - 6:   **for**  $1 \leq j \leq T_c$  **do**
  - 7:     Load  $\mathbf{K}_j, \mathbf{V}_j$  from HBM to on-chip SRAM.
  - 8:     On chip, compute  $\mathbf{S}_i^{(j)} = \mathbf{Q}_i \mathbf{K}_j^T \in \mathbb{R}^{B_r \times B_c}$ .
  - 9:     On chip, compute  $m_i^{(j)} = \max(m_i^{(j-1)}, \text{rowmax}(\mathbf{S}_i^{(j)})) \in \mathbb{R}^{B_r}, \tilde{\mathbf{P}}_i^{(j)} = \exp(\mathbf{S}_i^{(j)} - m_i^{(j)}) \in \mathbb{R}^{B_r \times B_c}$  (pointwise),  $\ell_i^{(j)} = e^{m_i^{(j-1)} - m_i^{(j)}} \ell_i^{(j-1)} + \text{rowsum}(\tilde{\mathbf{P}}_i^{(j)}) \in \mathbb{R}^{B_r}$ .
  - 10:    On chip, compute  $\mathbf{O}_i^{(j)} = \text{diag}(e^{m_i^{(j-1)} - m_i^{(j)}})^{-1} \mathbf{O}_i^{(j-1)} + \tilde{\mathbf{P}}_i^{(j)} \mathbf{V}_j$ .
  - 11:   **end for**
  - 12:   On chip, compute  $\mathbf{O}_i = \text{diag}(\ell_i^{(T_c)})^{-1} \mathbf{O}_i^{(T_c)}$ .
  - 13:   On chip, compute  $L_i = m_i^{(T_c)} + \log(\ell_i^{(T_c)})$ .
  - 14:   Write  $\mathbf{O}_i$  to HBM as the  $i$ -th block of  $\mathbf{O}$ .
  - 15:   Write  $L_i$  to HBM as the  $i$ -th block of  $L$ .
  - 16: **end for**
  - 17: Return the output  $\mathbf{O}$  and the logsumexp  $L$ .
-

# FlashAttention: 性能分析

- ❖ 考虑8B模型:  $\text{seq\_len}=8192$ ,  $\text{head\_dim}=128$
- ❖ 假设blocksize是 $128(Q) * 128(D) * 64(K)$
- ❖ 计算问题分为 $64*128$ 个tile
- ❖ 每个tile需要读取KV:  $2*64*128*2 = 0.03125\text{MB}$
- ❖ 但是每个tile要算两个矩阵乘:  $2*128*128*64*2 = 4\text{Mflops}$
- ❖ 计算访存比 =  $128 \text{ flops} / \text{Byte}$
- ❖ causal attention只需要算一半的tile, 共计128MB的IO (节省为1/8)
- ❖ 如果打满1008GB显存带宽, 实际能跑到139Tflops
- ❖ 实测能跑到170+Tflops (命中L2? + GPU turbo)

# 结构参数

| 模型版本    | 参数量  | 层数  | Hidden Dim | Heads | FFN Dim | Max Seq Len | Vocab Size |
|---------|------|-----|------------|-------|---------|-------------|------------|
| LLaMA-1 | 7B   | 32  | 4096       | 32    | 11008   | 2048        | 32000      |
| LLaMA-1 | 13B  | 40  | 5120       | 40    | 13824   | 2048        | 32000      |
| LLaMA-1 | 33B  | 60  | 6656       | 52    | 17920   | 2048        | 32000      |
| LLaMA-1 | 65B  | 80  | 8192       | 64    | 22016   | 2048        | 32000      |
| LLaMA-2 | 7B   | 32  | 4096       | 32    | 11008   | 4096        | 32000      |
| LLaMA-2 | 13B  | 40  | 5120       | 40    | 13824   | 4096        | 32000      |
| LLaMA-2 | 70B  | 80  | 8192       | 64    | 28672   | 4096        | 32000      |
| LLaMA-3 | 1B   | 24  | 2048       | 16    | 5504    | 2048        | 128256     |
| LLaMA-3 | 8B   | 32  | 4096       | 32    | 14336   | 8192        | 128256     |
| LLaMA-3 | 70B  | 80  | 8192       | 64    | 28672   | 8192        | 128256     |
| LLaMA-3 | 405B | 128 | 12288      | 96    | 49152   | 128K        | 128256     |

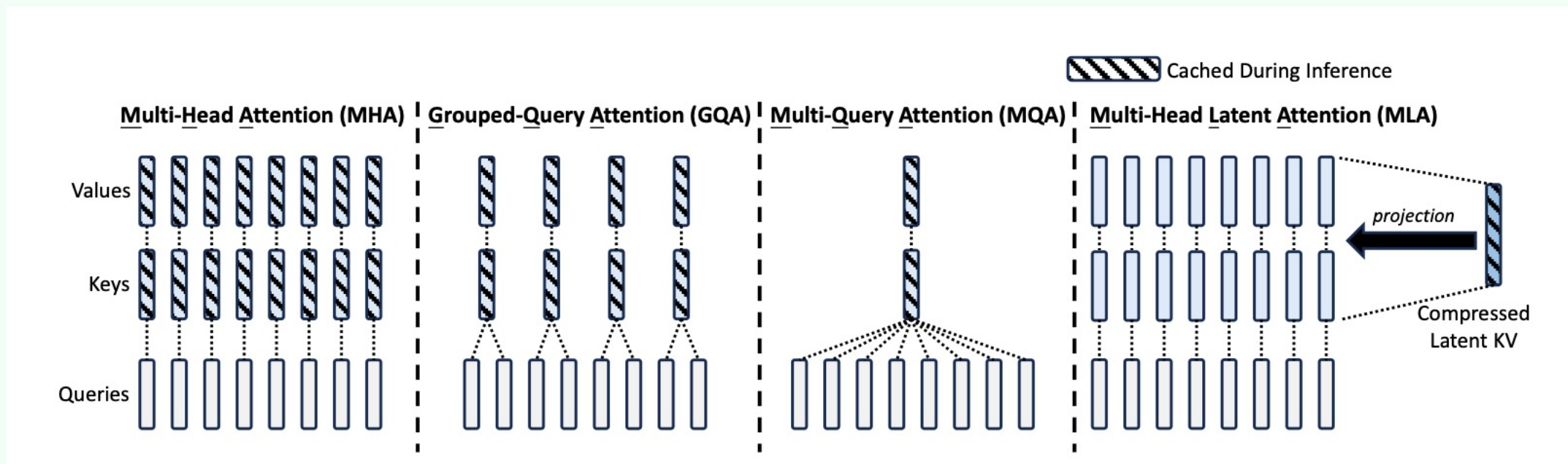
## MHA, MQA, GQA

问题: Kvcache太大。DeepSeekv3使用 $L=61$ ,  $\text{numheads}=128$ ,  $K\text{dim}=192$ ,  $V\text{dim}=128$

使用MHA需要约2.4MB / token

128K上下文需要298GB KVCache

思路: 减少Kvhead的数量



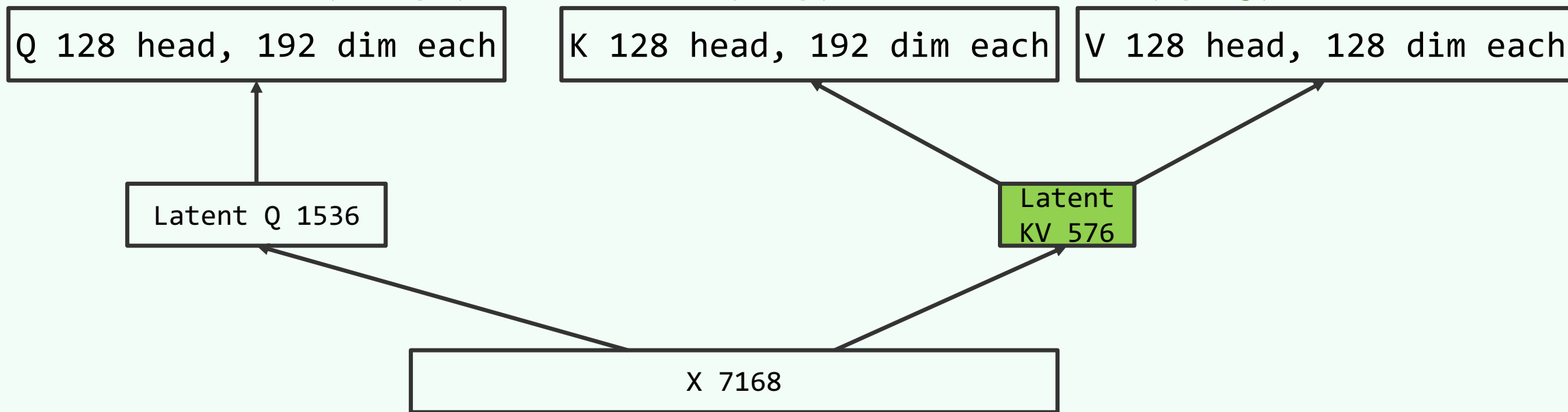
## Multi-latent attention

❖ 思路: 将attention projection变成MLP, 增加“压缩-解压”, 引入瓶颈层

❖ Kvcache:  $576*61 = 34\text{KB}$  / token, **71倍压缩率, 128K上下文仅需4.2GB!**

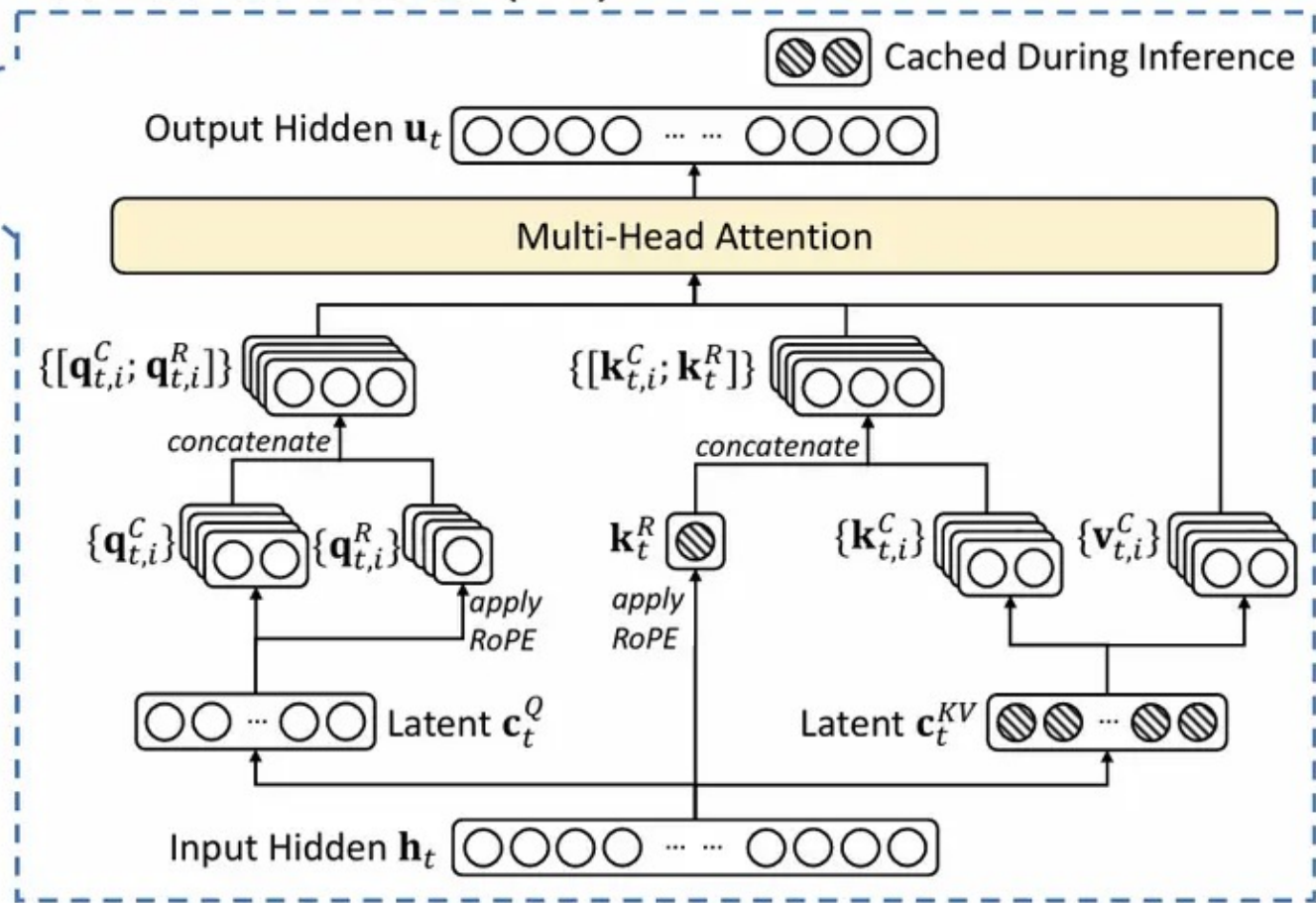
❖ 参数量:  $7168*(1536+576)$  (compress) +  $1536*128*192$  (Q up)

+  $7168*64(\text{KVrope})+576*128*256$  (KVup) +  $128*128*7168$  (Oproj) = 11.56B



$$\begin{aligned}
 \mathbf{c}_t^Q &= W^{DQ} \mathbf{h}_t, \\
 [\mathbf{q}_{t,1}^C; \mathbf{q}_{t,2}^C; \dots; \mathbf{q}_{t,n_h}^C] &= \mathbf{q}_t^C = W^{UQ} \mathbf{c}_t^Q, \\
 [\mathbf{q}_{t,1}^R; \mathbf{q}_{t,2}^R; \dots; \mathbf{q}_{t,n_h}^R] &= \mathbf{q}_t^R = \text{RoPE}(W^{QR} \mathbf{c}_t^Q), \\
 \mathbf{q}_{t,i} &= [\mathbf{q}_{t,i}^C; \mathbf{q}_{t,i}^R], \\
 \boxed{\mathbf{c}_t^{KV}} &= W^{DKV} \mathbf{h}_t, \\
 [\mathbf{k}_{t,1}^C; \mathbf{k}_{t,2}^C; \dots; \mathbf{k}_{t,n_h}^C] &= \mathbf{k}_t^C = W^{UK} \mathbf{c}_t^{KV}, \\
 \boxed{\mathbf{k}_t^R} &= \text{RoPE}(W^{KR} \mathbf{h}_t), \\
 \mathbf{k}_{t,i} &= [\mathbf{k}_{t,i}^C; \mathbf{k}_t^R], \\
 [\mathbf{v}_{t,1}^C; \mathbf{v}_{t,2}^C; \dots; \mathbf{v}_{t,n_h}^C] &= \mathbf{v}_t^C = W^{UV} \mathbf{c}_t^{KV}, \\
 \mathbf{o}_{t,i} &= \sum_{j=1}^t \text{Softmax}_j \left( \frac{\mathbf{q}_{t,i}^T \mathbf{k}_{j,i}}{\sqrt{d_h + d_h^R}} \right) \mathbf{v}_{j,i}^C, \\
 \mathbf{u}_t &= W^O [\mathbf{o}_{t,1}; \mathbf{o}_{t,2}; \dots; \mathbf{o}_{t,n_h}],
 \end{aligned}$$

## Multi-Head Latent Attention (MLA)



# MLA解码：计算QK

❖ 训练模式：直接暴力解压QKV并计算FlashAttention

❖ 推理模式：不显式解压，直接通过矩阵乘法完成解码计算

$$\begin{array}{c} \text{dim*512} \\ 1*1536 \quad 1536*\text{dim} \quad 512*\text{seqlen} \\ q_{t,i} K_i^T = (c_t^Q) (W_i^{UQ})^T W_i^{UK} (c^{KV})^T \\ 1*\text{dim} \quad \text{dim*seqlen} \\ = \ell_{t,i}^Q (c^{KV})^T \quad \begin{array}{l} \text{所有的head} \\ \text{共享, 只要} \\ \text{读一次} \end{array} \\ 1*512 \quad 512*\text{seqlen} \end{array}$$

凑成矩阵

$$\begin{array}{c} \text{\#heads} * \text{seqlen} \quad \text{\#heads} * 512 \\ S_t = \ell_t^Q (c^{KV})^T \\ 512 * \text{seqlen} \end{array}$$

t: query token id

j: kv token id

i: head id

- 因为latent KV是所有head共享的
- 所以一个矩阵乘法可以同时把所有head的S（矩阵）一起算出来
- 为什么MHA不能一起算？

MLA: ([heads, 1, dim] \* [1, seqlen, dim]).sum(2)

MHA: ([heads, 1, dim] \* [heads, seqlen, dim]).sum(2)

# MLA解码：计算PV

$$\diamond o_{t,i} = p_{t,i} V$$

1\*dim    1\*seqlen    seqlen\*dim

$$\diamond = p_{t,i} c^{KV} (W_i^{UV})^T$$

不分head  
seqlen\*512    512\*dim

结合律：先算 $V=c*W$ 还是先算 $r=p*c$ ??

$$\text{令 } R_{t,i} = (p_t c^{KV})_i$$

seqlen\*512  
#heads\*seqlen

$$\diamond o_{t,i} = R_{t,i} (W_i^{UV})^T$$

512\*128  
#heads\*512

同样是矩阵乘法！

## 总结: MLA

$$l_t^Q = (c_t^Q) (W_i^{UQ})^T W_i^{UK}$$

$$R_{t,i} = (p_t c^{KV})_i$$

$$S_t = \ell_t^Q (c^{KV})^T$$

$$o_{t,i} = R_{t,i} (W_i^{UV})^T$$

❖ 计算 $l$ :  $7168*1536+1536*192+192*512$

❖ 计算 $S$ :  $\text{seqlen}*\text{heads}*576$

❖ 计算 $R$ :  $\text{seqlen}*\text{heads}*576$

两个 $O(\text{seqlen})$ 的都是矩阵乘法

❖ 计算 $O$ :  $\text{heads}*\text{dim}*512$

❖ 访存: 读取 $W_i^{UQ}, W_i^{UK}, W_i^{UV}$ :  $2*(1536*128*128+512*128*128*2)=80\text{MB}$

❖ 访存: 两个 $\text{seqlen}*\text{heads}*512$ 的GEMM, 若 $\text{tile size}$ 是 $32, 128, 512$ 则

❖ 读 $l, c$ 、写 $R$ :  $128*512+\text{seqlen}*512+128*512$

## 总结: MLA

- ❖ 访存: 读取  $W_i^{UQ}, W_i^{UK}, W_i^{UV}$ :  $2 * (1536 * 128 * 128 + 512 * 128 * 128 * 2) = 80\text{MB}$
- ❖ 访存: 两个  $\text{seqlen} * \text{heads} * 512$  的 GEMM, 若  $\text{tilesize}$  是 32, 128, 512 则
- ❖ 读1,c、写R:  $2 * (128 * 512 + \text{seqlen} * 512 + 128 * 512) = 8.25\text{MB} (@8192) \quad 125.25\text{MB} (@120K)$
- ❖ 计算量:  $1.25\text{Gflops} (@8192) \quad 18.9\text{Gflops} (@128K)$
- ❖ 计算访存比:  $13.5\text{flops/Byte} (@8192, \text{bs}=1) \quad 87.8\text{flops/byte} (@128000, \text{bs}=1)$
- ❖ (理论上) 单卡4090读取一次Kvcache =  $61 * 125.25 / (1008 * 1024) = 7.5\text{ms}$
- ❖ 相比暴力解码后算MHA: 计算 =  $4 * \text{seqlen} * \text{numheads} * \text{headdim} = 0.5\text{Gflops} (@8192) / 8\text{Gflops} (@128K)$
- ❖ 访存: 读取整个KV =  $2 * 2 * \text{seqlen} * \text{numheads} * \text{headdim} = 512\text{MB} (@8192) / 8000\text{MB} (@120K)$
- ❖ 计算访存比 =  $1\text{flop} / \text{byte}$

DeepSeek真厉害!

## 反而比MHA更好

| Attention Mechanism           | KV Cache per Token (# Element)            | Capability |
|-------------------------------|-------------------------------------------|------------|
| Multi-Head Attention (MHA)    | $2n_h d_h l$                              | Strong     |
| Grouped-Query Attention (GQA) | $2n_g d_h l$                              | Moderate   |
| Multi-Query Attention (MQA)   | $2d_h l$                                  | Weak       |
| MLA (Ours)                    | $(d_c + d_h^R)l \approx \frac{9}{2}d_h l$ | Stronger   |

# 为什么(128,128)的KV能压缩进512维latent?

## ❖ 压缩感知理论, 有限等距性质 (Restricted Isometry Property)

对于任意高维稀疏向量 $x$ , 利用RIP矩阵 $\Phi$  (如randn)  
投影到低维空间后保范数。

$$(1 - \delta)\|x\|^2 \leq \|\Phi x\|^2 \leq (1 + \delta)\|x\|^2$$

- ❖ 将128\*128维空间中的128个128维超平面 (每个head的QK) 嵌入到一个512维的空间应该是简单的.....
- ❖ 每个head的QK是一个16384维向量, 但只有128维非零
- ❖ 区别: RIP是保距离, 我们需要的是保内积

# 为什么(128,128)的KV能压缩进512维latent?

❖ 根据RIP (Restricted Isometry Property) 理论

❖ 为了保住  $s$ -稀疏向量的内积, 所需的测量维度  $M$  (即 Latent Dim) 下界大约是:

$$M \approx C \cdot s \cdot \log\left(\frac{D}{s}\right)$$

❖  $s$  (稀疏度) =128,  $D$  (总维度) =16384

❖  $M \approx 896C$

❖ 核心:  $M$ 不需要是 $\text{numheads} \cdot \text{headdim}$ , 只需要是 $\text{headdim} * \log(\text{numheads})$ !

❖ 实际: 投影矩阵非随机, 而是学习的, 可能比该结果更优

**谢谢**

**清华大学计算机系 陈键飞**

**《大模型计算》课程团队**