

PR #49815 完整报告

vllm-project/vllm

[Bugfix][MiMo] Apply vision attention sinks in the window attention path

合并时间: 2026-08-11 06:36

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/49815>

执行摘要

- 一句话: 修复 MiMo 窗口注意力路径丢失视觉 sinks
- 推荐动作: 值得精读。两个设计决策很有参考价值: 一是用误差表量化区分两种 sinks 语义 (key 0 偏置 vs null logit), 避免把 `s_aux` 当作等价实现直接接入; 二是选择「内核内一步修正」而非「事后 LSE 修正」, 以可维护性为优先, 并配套性能数据 (2-4 倍加速) 证明取舍。对多模态模型接入自定义注意力语义、以及向 Triton 内核扩展可选模式的团队有直接借鉴意义。

功能与动机

Issue #47864 用户报告 `mimo_v2_omni.py` 视觉窗口注意力路径报错。根因在于 `self.sinks` 只在非 `fullatt_block_indexes` 的块上分配, 而这些块正是走 `_forward_window_attn` 的全部块; 该路径用 `flash_attn_varlen_func` 且从不读取 `self.sinks`。MiMo-V2.5 checkpoint 为除 `[0, 9, 18, 27]` 外的 24 个块都存了 `visual.blocks.N.attn.sinks`, 参考实现会把 `sinks[h]` 加到每条序列第一个 key 的 logit 上; 缺失后视觉编码器输出与参考误差高达 0.048, 属于明显的数值回归。

实现拆解

变更入口是 `MiMoVisionAttention._forward_window_attn`, 内核、模型、测试三处联动:

1. 根因定位 (`vllm/model_executor/models/mimo_v2_omni.py`): `self.sinks` 仅在 `use_sink=True` 时分配, 而 window 块 (非 `fullatt_block_indexes` 的 24 块) 全部走 `_forward_window_attn`; 原调用 `flash_attn_varlen_func` 未读 `self.sinks`, 确认是“权重加载了但从未消费”的数据契约断点。
2. 内核扩展 (`vllm/v1/attention/ops/triton_prefill_attention.py`): `_fwd_kernel` 新增编译期常量 `SINKS_BIAS_KEY0`, `context_attention_fwd` 新增默认 `False` 的 `sinks_bias_key0` 参数。默认路径保持原有 null logit 行为 (`m_i` 初始化为 sink、`l_i` 为 1.0); 新模式空初始化 `m_i/l_i`, 并在每个 key block 内用 `tl.where(mask & (pos_k == 0), qk + sink, qk)` 给 key 0 分数加偏置, 使 softmax 在单遍内对偏置后的分数归一化。
3. 模型路径切换 (`vllm/model_executor/models/mimo_v2_omni.py`): `_forward_window_attn` 从 `flash_attn_varlen_func` 改为 `context_attention_fwd`, 传入 `sliding_window_q/k=w`、`is_causal=False`, 并按 `self.tp_rank` 切片 `self.sinks` (TP 兼容) 后以 `sinks_bias_key0=True` 传入; `sinks` is None 时行为不变。

4. 测试配套 (tests/models/multimodal/test_mimo_v2_omni.py, 新增) : 用 fp32 稠密窗口 softmax 参考实现 `_reference` 做 ground truth, 参数化 MHA (num_kv_heads=4) 与 GQA (num_kv_heads=2), 序列长度 [5, 37] 分别短于 / 长于窗口 8; main 上误差 $1.2e-1/2.1e-1$, 修复后约 $2.3e-03$, 阈值 $1e-2$ 。
5. 兼容性验证: tests/kernels/attention/test_triton_prefill_attention.py 20 个既有用例全部通过, 确认未破坏原有 null logit 语义。

关键文件:

- vllm/model_executor/models/mimo_v2_omni.py (模块 视觉注意力; 类别 source; 类型 core-logic; 符号 MiMoVisionAttention, `_forward_window_attn`) : 修复主体: `_forward_window_attn` 从 `flash_attn_varlen_func` 切换为 Triton `context_attention_fwd`, 按 TP rank 切片 `self.sinks` 并传 `sinks_bias_key0=True`, 让窗口路径真正消费 checkpoint 中的 sink 权重。
- vllm/v1/attention/ops/triton_prefill_attention.py (模块 注意力内核; 类别 source; 类型 core-logic; 符号 `_fwd_kernel`, `context_attention_fwd`) : 内核支撑: 新增 `SINKS_BIAS_KEY0` 编译期常量与 `sinks_bias_key0` 运行时参数, 区分「key 0 偏置」与「null logit」两种 sink 语义, 是本修复能在一遍 softmax 内完成的关键。
- tests/models/multimodal/test_mimo_v2_omni.py (模块 模型测试; 类别 test; 类型 test-coverage; 符号 `vision_attn_env`, `_reference`, `test_window_attention_applies_sinks`) : 新增测试: 用 fp32 稠密窗口 softmax 参考实现对比 `_forward_window_attn`, 参数化 MHA/GQA, 验证 sink 偏置生效 (main 上 $1.2e-1/2.1e-1$ 的误差降到 $2.3e-03$), 是本次修复正确性的直接证据。

关键符号: `_forward_window_attn`, `context_attention_fwd`, `_fwd_kernel`, `test_window_attention_applies_sinks`, `_reference`

关键源码片段

vllm/model_executor/models/mimo_v2_omni.py

修复主体: `_forward_window_attn` 从 `flash_attn_varlen_func` 切换为 Triton `context_attention_fwd`, 按 TP rank 切片 `self.sinks` 并传 `sinks_bias_key0=True`, 让窗口路径真正消费 checkpoint 中的 sink 权重。

```
# vllm/model_executor/models/mimo_v2_omni.py —— 窗口注意力路径
def _forward_window_attn(self, q, k, v, cu_seqlens, max_seqlen):
    """窗口注意力: 把每头 sink logit 直接偏置到各序列 key 0 的分数上。
```

```
    参考实现 modeling_mimo_v2.py 在 softmax 前给 key 0 加 sinks[h],
    这里通过 Triton 内核的 sinks_bias_key0 模式单遍完成偏置与归一化,
    避免 flash attention 事后 LSE 修正带来的维护成本。
    """
```

```
    from vllm.v1.attention.ops.triton_prefill_attention import (
        context_attention_fwd,
    )
```

```
    w = self.visual_token_window_size
```

```

output = torch.empty_like(q)
# checkpoint 里的 sinks 是全量头数，需按 TP rank 切片取本卡的头
head_start = self.tp_rank * self.num_heads_per_partition
sinks = (
    self.sinks[head_start : head_start + self.num_heads_per_partition]
    if self.sinks is not None
    else None
)
context_attention_fwd(
    q, k, v, output,
    b_start_loc=cu_seqlens[:-1],
    b_seq_len=cu_seqlens[1:] - cu_seqlens[:-1],
    max_input_len=max_seqlen,
    is_causal=False,
    softmax_scale=self.scale,
    sliding_window_q=w,
    sliding_window_k=w,
    sinks=sinks,
    sinks_bias_key0=True,
)
return output

```

vllm/v1/attention/ops/triton_prefill_attention.py

内核支撑：新增 `SINKS_BIAS_KEY0` 编译期常量与 `sinks_bias_key0` 运行时参数，区分「key 0 偏置」与「null logit」两种 sink 语义，是本修复能在一遍 softmax 内完成的关键。

vllm/v1/attention/ops/triton_prefill_attention.py —— `_fwd_kernel` 中的 `sinks` 分支

if USE_SINKS:

载入当前头 (cur_head) 的 sink logit，并换算到 log2 域以匹配 exp2

sink = tl.load(Sinks + cur_head) * 1.4426950408889634

if SINKS_BIAS_KEY0:

MiMo 参考语义：sinks[h] 加到每条序列 key 0 的 logit 上，

因此 softmax 从常规空状态开始 (m_i 为 -inf、l_i 为 0)，

在下方循环内对偏置后的分数一次性完成归一化。

m_i = tl.zeros([BLOCK_M], dtype=tl.float32) - float("inf")

l_i = tl.zeros([BLOCK_M], dtype=tl.float32)

else:

原有 null logit 语义 (GPT-OSS)：sink 只是分母里的一项常数，

不会进入分子；m_i 与 l_i 直接以 sink 起步。

m_i = tl.full([BLOCK_M], sink, dtype=tl.float32)

l_i = tl.full([BLOCK_M], 1.0, dtype=tl.float32)

else:

m_i = tl.zeros([BLOCK_M], dtype=tl.float32) - float("inf")

l_i = tl.zeros([BLOCK_M], dtype=tl.float32)

每个 key block 内：先算分数并应用滑动窗口掩码，再按需偏置 key 0

qk = tl.dot(q, k)

qk = tl.where(mask, qk * sm_scale, -1.0e8)

if USE_SINKS and SINKS_BIAS_KEY0:

```
# 仅对掩码内且位置为 key 0 的分数加 sink，放大或缩小该 key 的权重
qk = tl.where(mask & (pos_k == 0), qk + sink, qk)
```

tests/models/multimodal/test_mimo_v2_omni.py

新增测试：用 fp32 稠密窗口 softmax 参考实现对比 `_forward_window_attn`，参数化 MHA/GQA，验证 sink 偏置生效（main 上 $1.2e-1/2.1e-1$ 的误差降到 $2.3e-03$ ），是本次修复正确性的直接证据。

```
# tests/models/multimodal/test_mimo_v2_omni.py —— 测试的 ground truth
```

```
def _reference(q, k, v, cu_seqlens, sinks, scale):
```

```
    """稠密 fp32 窗口 softmax 参考：sink 加到每个序列 key 0 的分数上。
```

```
    逐序列取 q/k/v 并转 fp32，GQA 时把 KV 头 repeat_interleave
    到与 Q 头一致，窗口外用 -inf 掩码，作为测试对比基准。
```

```
    """
```

```
    groups = q.shape[1] // k.shape[1]
```

```
    out = torch.empty_like(q, dtype=torch.float32)
```

```
    for start, end in zip(cu_seqlens[:-1].tolist(), cu_seqlens[1:].tolist()):
```

```
        qs = q[start:end].float()
```

```
        ks = k[start:end].float().repeat_interleave(groups, dim=1)
```

```
        vs = v[start:end].float().repeat_interleave(groups, dim=1)
```

```
        scores = torch.einsum("qhd,khd->hqk", qs, ks) * scale
```

```
        scores[..., 0] += sinks.float().view(-1, 1) # MiMo 语义：key 0 加偏置
```

```
        pos = torch.arange(end - start, device=q.device)
```

```
        outside = (pos.view(-1, 1) - pos.view(1, -1)).abs() > WINDOW
```

```
        scores.masked_fill_(outside, -torch.inf)
```

```
        out[start:end] = torch.einsum("hqk,khd->qhd", scores.softmax(-1), vs)
```

```
    return out
```

评论区精华

讨论围绕 `sinks` 的数学语义与实现路径展开，是典型的“用测量数据驱动设计决策”的案例：

- Isotr0py 先问「为什么不直接传 `s_aux=self.sinks`」，almogtavor 回应：FA2 会抛 `NotImplementedError`: FA2 does not support `s_aux`（只有 FA3/FA4 可用），且 `s_aux` 是 null logit 语义，只膨胀分母不进分子，与 MiMo 参考实现的 key 0 偏置不是同一运算——误差表显示 `s_aux` 路径误差 0.052，而本方案 0.0021。
- Isotr0py 建议「可以回退到 Triton 内核（pre-Hopper 设备）」；almogtavor 核查后确认现有 `USE_SINKS` 初始化（`m_i=sink`，`l_i=1.0`）与 `s_aux` 计算完全一致，同样不是 MiMo 语义，因此必须新增 `sinks_bias_key0` 模式，并指出窗口化路径的 $O(\text{seq_len}^2)$ 掩码问题已由 #50776 单独修复。
- 最终 Isotr0py 拍板：「I think using/correcting the triton kernel should be fine, LSE correction will increase the maintenance effort.」，放弃事后 LSE 修正路线，采纳内核内一步修正方案。
- `s_aux` 是否可直接使用（`sinks` 语义之争）（design）：放弃 `s_aux` / null logit 方案，采用 key 0 偏置语义。

- LSE 事后修正 vs Triton 内核内一步修正 (design): 采纳 Triton sinks_bias_key0 模式, LSE 修正方案被放弃。
- 现有 Triton USE_SINKS 语义验证与扩展 (correctness): 新增编译期常量 SINKS_BIAS_KEY0 与运行时参数 sinks_bias_key0, 默认行为不变, 既有 20 个内核测试通过。

风险与影响

- 风险:
 1. 路径级回归: 窗口注意力从 flash_attn_varlen_func 切到 Triton context_attention_fwd, 对所有 GPU 生效 (sink 只是其中一环)。Triton 内核的窗口化掩码原先有 $O(\text{seq_len}^2)$ 逐块遍历问题, 性能可行性依赖 #50776 的修复, 后者若回滚或行为变化会直接影响本路径。
 2. 平台兼容性: _forward_window_attn 现在依赖 vllm.v1 的 Triton 内核, 且测试 skipif 限定 CUDA, ROCm/XPU/CPU 等平台没有任何验证; 若这些平台的内核入口或数值行为不同, 可能引入新的平台差异。
 3. 数值正确性: 新偏置依赖 `tl.where(mask & (pos_k == 0), ...)`; 若滑动窗口掩码与 pos_k 位置判定在某个 key block 内不一致 (例如 block 跨序列边界), 可能漏加或错加偏置。测试仅覆盖 window=8 的小配置, 未覆盖 window 小于 key block 大小等边界。
 4. TP 切分: sinks[head_start:head_start + num_heads_per_partition] 依赖 self.tp_rank 语义正确, TP > 1 场景没有直接测试覆盖。- 影响: 用户侧: 修复 XiaomiMiMo/MiMo-V2.5 视觉编码器窗口注意力的数值错误 (与参考实现误差从 0.048 量级降到 bf16 噪声底 $2.1e-3$), 视觉 token 表征恢复与官方实现一致, 多模态推理质量回正。系统侧: context_attention_fwd 以默认 False 的参数向后兼容扩展, 既有 null logit (GPT-OSS) 语义不变; 为 Triton 注意力内核家族新增了「key 0 偏置」这一可复用语义, 未来其他模型可直接接入。团队侧: 新增的稠密参考对比测试给出了多模态注意力语义验证的模板; 但窗口路径从 flash-attn 迁到 vLLM 自持 Triton 内核后, 维护责任转移到自家内核。- 风险标记: 窗口注意力路径换用 Triton 内核, 依赖 #50776 的窗口掩码性能修复, 仅 CUDA 平台有测试覆盖, TP 分片逻辑无直接测试

关联脉络

- PR #50776 Triton sliding-window masking fix (讨论中提及, 标题未在上下文提供): PR 正文与 review 多次提到: Triton 内核原先逐 key block 遍历后再掩码, 窗口化路径代价 $O(\text{seq_len}^2)$, 该问题由 #50776 单独修复。本 PR 把窗口路径切到 Triton 内核, 性能可行性依赖此修复。
- PR #51734 replace batch_norm to numerically identical without cudnn: 同为多模态视觉模型推理正确性修复, 改动 vllm/model_executor/models/vision.py 并配套 tests/models/test_vision.py, 与本次 MiMo 视觉注意力修复同属多模态视觉模型质量收尾工作。
- PR #51461 [MM][CG][BugFix] Fix Ernie-4.5-VL encoder CG postprocess for multi-path outputs: 同为多模态视觉编码器 bugfix, 修复编码器路径在 CUDA graph / 多路径下的行为, 与本 PR 的视觉注意力路径修复属同一功能演进线。