

PR #50365 完整报告

vllm-project/vllm

[Perf][Sparse MLA] Drop the atomic contention in the index remap

合并时间: 2026-08-09 07:32

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/50365>

执行摘要

- 一句话: 消除稀疏 MLA 索引重映射原子竞争, kernel 提速最高 3.18x
- 推荐动作: 值得精读。这是一次教科书式的 kernel 竞争消除: 先定位“计数是 tile 间唯一通信原因”, 再用布局重构 (单 program 独占整行) 替代原子操作, 并配套 bit-exact 随机化验证与幂次边界测试。核心设计决策 (`_remap_tiling` 的幂次守卫、`SINGLE_TILE` 编译期分支、`torch.empty` 替代 `torch.zeros`) 都可直接借鉴。建议关注后续 AsariAI 与 vLLM 在稀疏 MLA 上的进一步协作。

功能与动机

稀疏 MLA 的 index remap 每行 2048 宽被切成 16 个列 tile, valid-count 路径上每个 tile 都向同一个 per-row 计数器 `atomic_add`, DCP 压缩路径更把该计数器当作原子槽位分配器——如 PR body 所述 'Counting is the only reason the tiles need to talk to each other'。该 kernel 每个 attention 层每个 decode step 执行一次, 16 路竞争随 batch 放大 (16 tile 与 4 tile 的差距随 batch 增长扩大, 是竞争的特征签名)。优化思路来自 AsariAI 博客 (<https://asari.ai/blog/inference-optimization>), 作者在 body 中明确邀请对方继续提交 PR。

实现拆解

1. 定位竞争源: `vllm/v1/attention/backends/mla/sparse_utils.py` 的 `_convert_req_index_to_global_index_kernel` 中, `COUNT_VALID` 路径所有 tile 向 `valid_count_ptr + token_id` 做 `tl.atomic_add`; `COMPACT_TO_FRONT` 路径把同一计数器当作槽位分配器, 16 tile 竞争同一地址。
2. 内核分支化: 新增编译期常量 `SINGLE_TILE`。为真时 `COMPACT_TO_FRONT` 的 base 直接取 0、计数用 `tl.store` 直写; 非压缩的 `COUNT_VALID` 路径同样以 `store` 替代 `atomic_add`。语义等价于“单 tile 的行内归约即行总数”。
3. 新增 `_remap_tiling` 布局决策: 仅在需要计数、且 `NUM_TOPK_TOKENS` 是 2 的幂 (`triton.next_power_of_2` 相等判断, 保证单个 `tl.arange` 可表达整行) 时选择单 tile 布局 (`block_n = NUM_TOPK_TOKENS`、`tiles_per_row = 1`、`num_warps = 8`); 否则回退原 128 宽、4 tile 布局 (`num_warps = 4`)。
4. 两个 host 入口统一接入: `triton_convert_req_index_to_global_index` 与 `triton_filter_and_convert_dcp_index` 都经 `_remap_tiling` 推导 `block_n`、`tiles_per_row`、`num_warps` 并透传 `SINGLE_TILE`; `valid_counts` 缓冲在单 tile 路径改用 `torch.empty` (零初始化只为 `atomic` 累加所需), 顺带省掉初始化开销。

5. 测试配套: tests/v1/attention/test_sparse_mla_backends.py 的 test_triton_convert_returns_valid_counts 参数化 [128, 384], tests/v1/attention/test_indexer_dcp_localize.py 的 test_dcp_filter_compaction_matches_reference 参数化 [1024, 384]; 384 特意覆盖非 2 的幂的 tiled atomic 回退路径。PR body 声称 1728 组随机化 A/B 配置 (token 数 1-1024、无效比例 0-100%、越界 block id、DCP size 2/4 与 interleave 1/64) 与旧实现逐位一致, 且压缩前缀顺序由“未指定”变为确定性 in-order。

关键文件:

- vllm/v1/attention/backends/mla/sparse_utils.py (模块 稀疏注意力; 类别 source; 类型 core-logic; 符号 _remap_tiling, _convert_req_index_to_global_index_kernel, triton_convert_req_index_to_global_index, triton_filter_and_convert_dcp_index) : 核心改动文件: 新增 SINGLE_TILE 编译期分支与 _remap_tiling 布局决策, 消除计数路径的跨 tile 原子竞争, 并将 valid_counts 缓冲从零初始化改为按路径选择 torch.empty/torch.zeros。
- tests/v1/attention/test_sparse_mla_backends.py (模块 后端测试; 类别 test; 类型 test-coverage; 符号 test_triton_convert_returns_valid_counts) : 将 test_triton_convert_returns_valid_counts 参数化为 [128, 384], 384 非 2 的幂, 用于覆盖 tiled atomic 回退路径的计数正确性。
- tests/v1/attention/test_indexer_dcp_localize.py (模块 索引器; 类别 test; 类型 test-coverage; 符号 test_dcp_filter_compaction_matches_reference) : 将 test_dcp_filter_compaction_matches_reference 参数化为 [1024, 384], 验证单 tile 与多 tile 两种压缩布局下, 内核内压缩结果均与参考 filter + sort/gather 实现集合一致。

关键符号: _remap_tiling, _convert_req_index_to_global_index_kernel, triton_convert_req_index_to_global_index, triton_filter_and_convert_dcp_index, test_triton_convert_returns_valid_counts, test_dcp_filter_compaction_matches_reference

关键源码片段

vllm/v1/attention/backends/mla/sparse_utils.py

核心改动文件: 新增 SINGLE_TILE 编译期分支与 _remap_tiling 布局决策, 消除计数路径的跨 tile 原子竞争, 并将 valid_counts 缓冲从零初始化改为按路径选择 torch.empty/torch.zeros。

```
# vllm/v1/attention/backends/mla/sparse_utils.py
# 稀疏 MLA 索引重映射 kernel 的计数与压缩路径。
# 旧实现把每行切成 16 个 128 宽的列 tile, 每个 tile 都向同一个 per-row
# 计数器 atomic_add, decode 步骤里形成 16 路写竞争; DCP 压缩路径还把
# 该计数器当作原子槽位分配器, 竞争更严重。
```

```
@triton.jit
```

```
def _convert_req_index_to_global_index_kernel(
    req_id_ptr, # int32 [num_tokens]
    block_table_ptr, # int32 [num_requests, max_num_blocks_per_req]
    token_indices_ptr, # int32 [num_tokens, NUM_TOPK_TOKENS]
```

```

out_ptr, # int32 [num_tokens, NUM_TOPK_TOKENS]
valid_count_ptr, # int32 [num_tokens], 每行有效计数输出
...
COUNT_VALID: tl.constexpr,
# 新增 SINGLE_TILE: BLOCK_N 与 NUM_TOPK_TOKENS 相等时, 单个 program
# 独占整行, 每行计数就是它自己的寄存器归约, 无需与其它 tile 通信。
SINGLE_TILE: tl.constexpr,
COMPACT_TO_FRONT: tl.constexpr,
...
):
    token_id = tl.program_id(0)
    tile_id = tl.program_id(1)
    # 每个 program 覆盖 BLOCK_N 个连续列; 单 tile 时即整行
    indice_id = tile_id * BLOCK_N + tl.arange(0, BLOCK_N)
    # 行加载、DCP 去交织、block_table 查找与 -1 兜底逻辑未变, 从略

    if COMPACT_TO_FRONT:
        # DCP 过滤路径: 把本 rank 拥有的槽位压到行首连续前缀
        # [0, valid_count), 行尾保持 -1 (out 缓冲已预填充 -1)。
        is_valid = (~is_invalid_tok).to(tl.int32)
        # 独占前缀和给出每个有效 lane 在行内的偏移
        local_offset = tl.cumsum(is_valid) - is_valid
        tile_valid_count = tl.sum(is_valid)
        if SINGLE_TILE:
            # 没有竞争 tile, 压缩基址静态为 0:
            # 原子槽位分配器退化为一次普通 store
            base = 0
            tl.store(valid_count_ptr + token_id, tile_valid_count)
        else:
            # 多 tile 时一次 atomic_add 预留连续基址
            base = tl.atomic_add(valid_count_ptr + token_id, tile_valid_count)
            dest = base + local_offset
            out_ptr_dest = out_ptr + token_id * out_stride0 + dest * out_stride1
            tl.store(out_ptr_dest, out_val, mask=is_valid == 1)
    else:
        # 原位输出路径 (输入列与输出列一致)
        out_ptr_ij = out_ptr + token_id * out_stride0 + indice_id * out_stride1
        tl.store(out_ptr_ij, out_val)
        if COUNT_VALID:
            tile_valid_count = tl.sum((~is_invalid_tok).to(tl.int32))
            if SINGLE_TILE:
                # 单 tile 的归约结果就是行总数, 直接 store
                tl.store(valid_count_ptr + token_id, tile_valid_count)
            else:
                tl.atomic_add(valid_count_ptr + token_id, tile_valid_count)

def _remap_tiling(
    NUM_TOPK_TOKENS: int, BLOCK_N: int, count_valid: bool

```

) -> tuple[bool, int, int, int]:

```
"""选择索引重映射 kernel 的列分块布局，返回 (single_tile, block_n,
tiles_per_row, num_warps)。
```

统计有效槽位是列 tile 之间唯一需要通信的地方，因此计数路径让单个 program 独占整行：计数退化为寄存器归约加普通 store，免掉 atomic 与计数器零初始化。整行是单个 tl.arange，要求行宽为 2 的幂；非 2 的幂宽度（如 384）继续走分块 atomic 累加路径。

```
"""
```

```
single_tile = (
    count_valid and triton.next_power_of_2(NUM_TOPK_TOKENS) == NUM_TOPK_TOKENS
)
if single_tile:
    return True, NUM_TOPK_TOKENS, 1, 8
return False, BLOCK_N, NUM_TOPK_TOKENS // BLOCK_N, 4
```

评论区精华

核心交锋只有一条但非常关键：yewentao256 在初始 diff (`single_tile = return_valid_counts` 无条件开启单 tile) 上提问“Will 384 be good for this change?”，直指 384 这类非 2 的幂行宽无法用单个 `tl.arange` 表达的问题；njhill 回复“Thanks @yewentao256 good catch :) I have pushed another commit”，追加提交 c351ef8 用 `triton.next_power_of_2` 守卫单 tile 生效条件，并给两个测试补上 384 参数化用例。该问题在合并前被捕获并固化进回归测试，体现了 review 的价值。

- 非 2 的幂 index_topk 在单 tile 路径下的正确性 (correctness): njhill 承认遗漏并追加提交 c351ef8 “Keep the tiled path for non-power-of-two index_topk”，用 `triton.next_power_of_2(NUM_TOPK_TOKENS) == NUM_TOPK_TOKENS` 限制单 tile 生效条件，并给两个测试补上 384 参数化用例覆盖 tiled atomic 回退路径。

风险与影响

- 风险：（1）行为契约风险：计数从多 tile 原子累加改为单 program 直写，隐含“每行恰好一个 program 写计数”的假设，SINGLE_TILE 目前只由 `_remap_tiling` 推导，未来新增非 2 的幂或超大行宽配置时需同步确认约束。（2）寄存器压力：单 tile 以 8 warps 处理 2048 列，寄存器 / 共享内存占用上升，极端配置下有 spilling 风险；GB200 上收益远超开销，但 ROCm 等其它架构未单独验证。（3）DCP 压缩前缀顺序从未指定变为确定性 in-order，属语义改善，但理论上需排查下游是否隐式依赖旧乱序。（4）影响面：flashmla、flashinfer、flashattn、xpu、rocm_aiter 共用该路径，回归影响半径大，依赖 1728 组随机化 A/B 验证与 384 参数化测试兜底。
- 影响：对用户：DeepSeek 类稀疏 MLA 模型 decode 吞吐与每 step 延迟直接受益，78 层 index_topk=2048 模型每 step 省约 93us (1 token) 至 270us (256 token)。对系统：改动局限在 `vllm/v1/attention/backends/mla/sparse_utils.py` 一个源文件，布局决策集中到 `_remap_tiling` 一处，后续新增 topk 宽度只需在该函数补策略。对团队：PR 展示了“先做竞争签名分析（16 tile 与 4 tile 差距随 batch 放大）再做结构改造”的优化方法论，以及用大规模随机化 A/B 验证 bit-exact 的测试纪律，可作为 kernel 优化 PR 的模板。

- 风险标记: 热点 decode 路径 kernel 变更, 多后端共享路径 (flashmla/flashinfer/flashattn/xpu/rocm_aiter), 单 tile 布局受 2 的幂限制, 计数器初始化语义变化 (zeros -> empty), 寄存器压力与 spilling 风险

关联脉络

- PR #51298 [DSv32/GLM Perf] Skip short prefill topk for dense mha layer, 97.9% kernel level latency reduction: 同为 DeepSeek V3.2 稀疏 MLA 注意力路径的 kernel 层性能优化, 与本 PR 处于同一 attention 代码域与优化节奏。
- PR #51425 [Perf] Narrow DeepSeek V3.2 eager CUDA graph region: 持续削减 DSv3.2 decode 图内 kernel 耗时, 与本 PR 的 remap 优化同属 decode 性能攻坚线。
- PR #51434 [Perf] Optimize DeepSeek V3.2 sequence parallelism: DSv3.2 序列并行全链路化, 稀疏 MLA 数据流随之演进, 可能影响 remap 的 DCP 语义边界。