

# PR #51841 完整报告

vllm-project/vllm

Avoid long-blocking H2D copies in ViT

合并时间: 2026-08-12 22:54

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/51841>

## 执行摘要

- 一句话: 消除 ViT 两处长阻塞 H2D 拷贝, 恢复异步流水
- 推荐动作: 值得精读。虽然改动仅 23 行、语义不变, 但 PR 的价值在于: 一是对 PyTorch H2D 内部行为 (strided 源 → pageable gather → 静默同步) 给出清晰机制解释; 二是 profiling + 微基准的定位方法规范; 三是 njhill 对 '避免多余拷贝' 的设计敏感度值得学习。建议关注后续是否有类似 '异步被静默破坏' 的排查需求, 可将此 PR 作为模板; 同时留意 Qwen3-VL 侧 pinned 缓冲区是否有必要做缓存复用。

## 功能与动机

PR body 明确指出: 'Two host-to-device copies on the per-iteration critical path can take a very long time to return. Despite `non_blocking=True`, the `cudaMemcpyAsync` call blocks the calling thread, and the time it takes tracks the amount of GPU work that has already been queued. That undoes the host run-ahead that asynchronous scheduling and CUDA graphs exist to build up.' 两个调用点从代码表面看都不可疑, 但 profiling 显示它们耗时随 GPU 已排队工作量的深度增长, 说明异步语义被静默破坏。

## 实现拆解

1. 定位阶段: 通过 profiling 在 V1 `gpu_model_runner.py::_prepare_inputs()` 的 M-RoPE 拷贝和 `qwen3_vl.py::rot_pos_emb()` 的视觉位置 id 拷贝两处发现超长 `cudaMemcpyAsync`, 并识别出共同模式——H2D 源数据要么是 strided view, 要么是 pageable 内存。
2. M-RoPE 修复 (`vllm/v1/worker/gpu_model_runner.py`): `mrope_positions` 被有意分配为 `[3, max_num_tokens + 1]` (末尾 dummy 列用于对 `torch.compile` 保持 non-contiguous), 因此 `cpu[:, :N]` 是 strided view; `copy_()` 无法把 strided 源表达为单次 `cudaMemcpyAsync`, 会先 gather 到连续的 pageable 临时缓冲区, 而 pageable H2D 会忽略 `non_blocking=True` 并在传输开始前同步。修复改为对 3 行分别调用 `copy_()`, 每一行在同一个 pinned 内存分配内连续, 从而真正走异步 pinned 路径。
3. Qwen3-VL 修复 (`vllm/model_executor/models/qwen3_vl.py`): `rot_pos_ids()` 从 numpy 构造 per-image 张量, 拼接后的 `pos_ids` 位于普通 pageable 主机内存, 直接 `.to(device, non_blocking=True)` 会同步。修复先用 `torch.empty(..., pin_memory=True)` 开辟 pinned 缓冲区 (尺寸由 `num_pos = sum(p.shape[0] for p in pos_ids)` 计算), 再用 `torch.cat(..., out=pinned)` 把拼接结果直接写入该缓冲区, 避免先拼到 pageable 内存再

pin\_memory() 带来的额外 host 侧拷贝。

4. 协同优化: njhill 的 review 建议被采纳, 将最初的链式 .pin\_memory() 改为上述 cat(out=...) 写法, 减少一次额外 host 拷贝; 第三、第四个 commit 依次补齐了 Co-authored-by 署名和 ruff-format 格式修正。
5. 测试与验证: 未新增测试文件, 依赖现有 Qwen-VL correctness 测试; PR body 提供 Nsight 前后对比截图证明两处长阻塞拷贝消失, 并声明输出与之前一致。Buildkite CI (#83439) 通过后由 Isotr0py approve 合入。

关键文件:

- vllm/v1/worker/gpu\_model\_runner.py (模块 模型执行; 类别 source; 类型 core-logic; 符号 \_prepare\_inputs): V1 模型执行热路径上的 M-RoPE 位置拷贝是本次改动的主战场: 将一次 strided 整体拷贝拆成 3 次逐行连续拷贝, 消除因 strided view 转 pageable gather 导致的隐式 stream 同步, 直接影响所有使用 M-RoPE 的模型 (Qwen2-VL、Qwen2.5-VL、Qwen3-VL) 每步迭代的 host run-ahead。
- vllm/model\_executor/models/qwen3\_vl.py (模块 模型实现; 类别 source; 类型 core-logic; 符号 rot\_pos\_emb): Qwen3-VL 视觉位置 id 的 H2D 拷贝是第二个阻塞点: rot\_pos\_ids() 从 numpy 构造张量导致拼接产物位于 pageable 内存, 修复改为将 cat 结果直接写入 pinned 缓冲区 (采用 njhill review 建议的 out=pinned 写法, 避免额外 host 拷贝), 使 H2D 回到异步路径。

关键符号: \_prepare\_inputs, rot\_pos\_emb

## 关键源码片段

### vllm/v1/worker/gpu\_model\_runner.py

V1 模型执行热路径上的 M-RoPE 位置拷贝是本次改动的主战场: 将一次 strided 整体拷贝拆成 3 次逐行连续拷贝, 消除因 strided view 转 pageable gather 导致的隐式 stream 同步, 直接影响所有使用 M-RoPE 的模型 (Qwen2-VL、Qwen2.5-VL、Qwen3-VL) 每步迭代的 host run-ahead。

```
# vllm/v1/worker/gpu_model_runner.py :: _prepare_inputs()
if self.uses_mrope:
    # 仅对使用 M-RoPE 的模型生效 (例如 Qwen2-VL、Qwen3-VL)。
    # mrope_positions 被分配为 [3, max_num_tokens + 1], 末尾的 dummy 列
    # 有意保留, 用来让 tensor 对 torch.compile 保持 non-contiguous,
    # 因此 cpu[:, :N] 是一个 strided view。
    # copy_() 无法把 strided 源表达为单次 cudaMemcpyAsync, 会先 gather
    # 到连续的 *pageable* 临时缓冲区; 而 pageable 的 H2D 会忽略
    # non_blocking=True, 并在传输开始前同步当前 stream。
    # 每个 row 在同一个 pinned 分配内是连续的, 所以逐行拷贝
    # 可以一直走 pinned 异步路径, 真正避免同步。
    for row in range(self.mrope_positions.gpu.shape[0]):
        self.mrope_positions.gpu[row, :total_num_scheduled_tokens].copy_(
            self.mrope_positions.cpu[row, :total_num_scheduled_tokens],
            non_blocking=True,
        )
```

```

elif self.uses_xdrope_dim > 0:
    # XD-RoPE (例如 HunYuan-VL) 保持原有整体拷贝逻辑。
    self.xdrope_positions.gpu[:, :total_num_scheduled_tokens].copy_(
        self.xdrope_positions.cpu[:, :total_num_scheduled_tokens],
        non_blocking=True,
    )

```

## vllm/model\_executor/models/qwen3\_vl.py

Qwen3-VL 视觉位置 id 的 H2D 拷贝是第二个阻塞点：rot\_pos\_ids() 从 numpy 构造张量导致拼接产物位于 pageable 内存，修复改为将 cat 结果直接写入 pinned 缓冲区（采用 njhill review 建议的 out=pinned 写法，避免额外 host 拷贝），使 H2D 回到异步路径。

```

# vllm/model_executor/models/qwen3_vl.py :: Qwen3VLImageEmbedding.rot_pos_emb()
def rot_pos_emb(self, grid_thw: list[list[int]]):
    max_grid_size = max(max(h, w) for _, h, w in grid_thw)
    pos_ids = [
        self.rot_pos_ids(h, w, self.spatial_merge_size)
        if t == 1
        else self.rot_pos_ids(h, w, self.spatial_merge_size).repeat(t, 1)
        for t, h, w in grid_thw
    ]
    # rot_pos_ids() 内部从 numpy 构造张量，拼接产物位于普通 pageable 内存；
    # 直接 .to(device, non_blocking=True) 会因 pageable 源而忽略 non_blocking
    # 并在传输前同步 stream。
    # 这里把拼接输出直接写入 pin_memory=True 的缓冲区（cat 的 out= 参数），
    # 既避免先拼到 pageable 内存再 pin_memory() 的额外 host 拷贝，
    # 又保证 H2D 源是 pinned 且连续的，从而进入真正的异步路径。
    num_pos = sum(p.shape[0] for p in pos_ids)
    pinned = torch.empty(
        (num_pos, pos_ids[0].shape[1]),
        dtype=pos_ids[0].dtype,
        pin_memory=True,
    )
    pos_ids = torch.cat(pos_ids, dim=0, out=pinned).to(
        self.device, non_blocking=True
    )

    # 使用 RotaryEmbedding 预计算的 cos_sin_cache
    cos, sin = self.rotary_pos_emb.get_cos_sin(max_grid_size)

    cos_combined = cos[pos_ids].flatten(1)
    sin_combined = sin[pos_ids].flatten(1)

    return cos_combined, sin_combined

```

## 评论区精华

1. Isotr0py 对 V2 的质疑 (vllm/v1/worker/gpu\_model\_runner.py) : 评论 '~~Do we have similar issues in v2?~~ NVM, I forgot we use UVA in V2', 并引用

vllm/v1/worker/gpu/mm/rope.py L98-L106 说明 Model Runner V2 通过 UVA 缓冲暂存位置信息，不执行此类 H2D 拷贝，因此不受影响。

2. njhill 对实现方式的优化 (vllm/model\_executor/models/qwen3\_vl.py)：指出链式 .pin\_memory() 会引入额外 host 拷贝，建议 'can avoid an extra host-side copy here by catting into an already-pinned tensor'，给出 torch.empty(pin\_memory=True) + torch.cat(..., out=pinned) 的具体代码；该建议被采纳，对应 commit 带有 Co-authored-by: Nick Hill。
3. 最终结论: Isotr0py APPROVED 并留言 'LGTM, thanks!', 无遗留未解决讨论。
  - Model Runner V2 是否存在相同问题 (question): V2 使用 UVA 方案，不受该问题影响，无需同步修改。
  - 用 cat(out=pinned) 消除额外 host 拷贝 (performance): 建议被采纳，随后的 commit 应用该写法并带上 Co-authored-by: Nick Hill。

## 风险与影响

- 风险:
  1. 热路径变更 (gpu\_model\_runner.py)：M-RoPE 分支位于 V1 每步迭代关键路径，由一次整体 copy\_() 变为 3 次逐行 copy\_()，增加了 2 次调用指令开销；普适情况下收益（消除长阻塞）远大于开销，但对 total\_num\_scheduled\_tokens 很小的浅队列场景收益可能有限，PR 未提供该场景的端到端基准。
  2. pinned 内存分配 (qwen3\_vl.py)：rot\_pos\_emb() 每次被调用都会新分配一块 pinned 内存 (cudaHostAlloc 本身可能较慢且不可换页)，对并发图像请求多的场景可能带来额外分配开销，代码未做缓存复用。
  3. 测试覆盖：没有新增任何测试文件，回归兜底仅靠现有 Qwen-VL correctness 测试；若未来有人改动 mrope\_positions 的内存布局（如去掉 dummy 列），逐行拷贝逻辑与注释中的前提会失去约束。
  4. 根因未完全确立：PR body 明确声明 'The root cause is not established'，阈值行为只来自微基准观察，未在真实模型配置端到端确认；若底层是 driver 或其他机制，其他 strided/pinned 组合仍可能出现类似静默同步。- 影响：用户侧：Qwen2-VL、Qwen2.5-VL、Qwen3-VL 等使用 M-RoPE 的视觉模型在 V1 引擎下每步迭代的延迟改善，尤其在 GPU 队列深、host run-ahead 被破坏的场景下效果显著；Qwen3-VL 视觉位置编码路径的同步消失。系统侧：V1 模型执行路径的 H2D 异步性恢复，与异步调度和 CUDA graphs 的设计目标一致；Model Runner V2 通过 UVA 方案天然规避此类问题，不受影响。团队侧：提供了一套可复用的排查范式——在 profiling 中识别超长 cudaMemcpyAsync、检查源缓冲区是否 strided 或 pageable、优先保证 pinned 且连续，对后续同类性能问题的定位有直接借鉴价值。- 风险标记：核心路径热路径变更，缺少新增测试覆盖，根因未完全确立，新增 pinned 内存分配

## 关联脉络

- PR #51738 [Perf] Avoid more GPU<->CPU syncs on the model execution path: 同属 V1 模型执行路径的性能优化线，目标是消除 GPU-CPU 同步、恢复异步执行，与本次 H2D 隐式同步修复关注点一致。

- PR #51913 [Attention] Move context\_lens\_tensor compute into GDN prefill path: 同为 V1 路径减少冗余计算与同步的性能改动，可用于对照 V1 性能优化节奏。
- PR #51843 [Bugfix] Disable fine-grained prefix-cache hits for incompatible hybrid KV layouts: 同属 V1 调度 / 执行路径的近期 bugfix，涉及相同 scheduler 与 model runner 相关模块，体现 V1 路径的持续演进。