

PR #44944 完整报告

vllm-project/vllm

[PERF] Fuse multi-group block table staged writes

合并时间: 2026-06-17 01:53

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/44944>

执行摘要

- 一句话: 融合多 KV-group 的 block table 分阶段写入为一个 Triton 内核
- 推荐动作: 值得精读设计模式: 如何将多个 StagedWriteTensor 的 staged writes 融合为一个内核, 保持单组路径不变。但性能收益需在自有环境复现后再决定是否启用。建议增加更全面的压力测试和跨后端验证。

功能与动机

PR body 明确指出目的是融合多 KV-group 的分阶段写入内核, 从而提升性能。作者在 H800 上测试显示高并发时 TPOT 有明显改善, 但 reviewer njhill 无法复现同等幅度的提升, 认为属于微优化。最终在作者坚持和 reviewer 批准下合并。

实现拆解

1. 修改 Triton 内核 `_apply_write_kernel`: 增加 `MULTI_GROUP` 编译常量和 `write_group_ids_ptr` 参数。当 `MULTI_GROUP=True` 时, 每个线程块先通过 `write_group_ids_ptr` 加载对应的 `group_id`, 然后从 `output_ptr` 和 `output_stride` (此时视为指针数组) 中加载该 `group` 的基地址和步长, 实现单个内核写多个 block table。
2. 新增 `FusedStagedWriter` 类 (`buffer_utils.py`): 维护 `group_ids`、`indices`、`starts`、`cu_lens` 四个 `UvaBufferPool`, 并在 `apply()` 方法中遍历所有 `StagedWriteTensor`, 收集其 staged write 内容并合并为一个批次调用 `_apply_write_kernel`, 最后清空所有 tensor 的 staged writes。
3. 修改 `BlockTables` 初始化 (`block_table.py`): 在 `__init__` 中根据 `num_kv_cache_groups > 1` 有条件地创建 `FusedStagedWriter` 实例, 分配足够容纳 `num_kv_cache_groups * max_num_reqs` 次写入的缓冲区。
4. 修改 `apply_staged_writes` 方法: 当 `num_kv_cache_groups == 1` 时沿用原单组路径; 多组时调用 `self.fused_writer.apply(self.block_tables, self.block_table_ptrs, self.block_table_strides)`, 并统一执行 `self.num_blocks.copy_to_uva()`。
5. 新增 CUDA 测试文件 (`test_gpu_block_table.py`): 包含 `test_block_tables_apply_staged_writes_fuses_kv_groups` (通过 monkeypatch 禁止单组 `apply_write` 被调用, 验证融合路径正确性, 包括 block ID 扩展、overwrite 与 append 模式) 和 `test_block_tables_apply_staged_writes_single_group` (验证单组路径仍走原有实现并正确写入)。

关键文件：

- vllm/v1/worker/gpu/buffer_utils.py (模块 缓冲工具；类别 source；类型 core-logic；符号 FusedStagedWriter, init, apply, _load_ptr)：核心变更文件：新增 FusedStagedWriter 类并修改 _apply_write_kernel 以支持 MULTI_GROUP 模式，实现多 block table 的融合写入。
- vllm/v1/worker/gpu/block_table.py (模块 块表；类别 source；类型 core-logic；符号 _load_ptr)：调用入口：修改 BlockTables 的初始化与 apply_staged_writes 方法，使其有条件地使用 FusedStagedWriter。
- tests/v1/worker/test_gpu_block_table.py (模块 块表测试；类别 test；类型 test-coverage；符号 test_block_tables_apply_staged_writes_fuses_kv_groups, fail_if_apply_write_called, test_block_tables_apply_staged_writes_single_group)：新增测试文件，覆盖融合路径和单组路径的正确性，包括 block ID 扩展、overwrite/append 模式以及清空检查。

关键符号：FusedStagedWriter.init, FusedStagedWriter.apply, BlockTables.apply_staged_writes, BlockTables.init, StagedWriteTensor.apply_write, _apply_write_kernel

关键源码片段

vllm/v1/worker/gpu/buffer_utils.py

核心变更文件：新增 FusedStagedWriter 类并修改 _apply_write_kernel 以支持 MULTI_GROUP 模式，实现多 block table 的融合写入。

FusedStagedWriter 类的 apply 方法：将多个 StagedWriteTensor 的 staged writes 合并为一个内核调用

```
class FusedStagedWriter:
```

```
    def apply(
        self,
        tensors: Sequence[StagedWriteTensor],
        output_ptrs: torch.Tensor, # 每个 group 的 block table GPU 数据指针
        output_strides: torch.Tensor, # 每个 group 的 stride(0)
```

```
) -> None:
```

```
    group_ids: list[int] = []
    indices: list[int] = []
    starts: list[int] = []
    contents: list[int | float] = []
    cu_lens: list[int] = []
```

```
    for group_id, t in enumerate(tensors):
```

```
        n = len(t._staged_write_indices)
```

```
        if n == 0:
```

```
            continue
```

```
        # 记录每个 write 对应的 group_id
```

```
        group_ids.extend([group_id] * n)
```

```
        indices.extend(t._staged_write_indices)
```

```
        starts.extend(t._staged_write_starts)
```

```

        content_base = len(contents)
        contents.extend(t._staged_write_contents)
        # cu_lens 需要偏移到全局 contents 中的位置
        cu_lens.extend(content_base + cu_len for cu_len in t._staged_write_cu_lens)

    if not group_ids:
        return

    # 将收集的数据拷贝到 GPU (UVA)
    group_ids_uva = self.group_ids.copy_to_uva(group_ids)
    indices_uva = self.indices.copy_to_uva(indices)
    starts_uva = self.starts.copy_to_uva(starts)
    cu_lens_uva = self.cu_lens.copy_to_uva(cu_lens)
    contents_gpu = async_tensor_h2d(contents, torch.int32, self.device)

    # 启动一次 Triton 内核处理所有 writes
    _apply_write_kernel([(len(group_ids),)](
        output_ptrs,
        output_strides,
        indices_uva,
        starts_uva,
        contents_gpu,
        cu_lens_uva,
        group_ids_uva, # 新增参数, 指示每个 write 属于哪个 group
        BLOCK_SIZE=1024,
        MULTI_GROUP=True, # 编译常量, 使 kernel 走多组分支
    ))
    # 清空所有 tensor 的 staged writes
    for t in tensors:
        t.clear_staged_writes()

# Triton 内核核心分支 (MULTI_GROUP=True 时)
# 每个 pid 处理一个 write: 先加载 group_id, 再通过 pointer-to-pointer 获取该 group 的 base_ptr
if MULTI_GROUP:
    group_id = tl.load(write_group_ids_ptr + pid)
    row_ptr = _load_ptr(output_ptr + group_id, tl.int32) # 双指针解引用
    row_stride = tl.load(output_stride + group_id)
else:
    row_ptr = output_ptr
    row_stride = output_stride

```

vllm/v1/worker/gpu/block_table.py

调用入口: 修改 BlockTables 的初始化与 apply_staged_writes 方法, 使其有条件地使用 FusedStagedWriter。

```

# BlockTables.__init__ 中新增: 仅当多组时创建 FusedStagedWriter
self.fused_writer: FusedStagedWriter | None = None
if self.num_kv_cache_groups > 1:
    self.fused_writer = FusedStagedWriter(

```

```

        self.device, self.num_kv_cache_groups * self.max_num_reqs
    )

# apply_staged_writes 方法：根据 group 数选择路径
def apply_staged_writes(self) -> None:
    if self.num_kv_cache_groups == 1:
        # 单组：仍使用原有的 apply_write 方法（单 kernel）
        self.block_tables[0].apply_write()
    else:
        # 多组：使用融合 writer，一次内核调用处理所有 group 的 staged writes
        assert self.fused_writer is not None
        self.fused_writer.apply(
            self.block_tables, self.block_table_ptrs, self.block_table_strides
        )
    # 同步 num_blocks 到 GPU
    self.num_blocks.copy_to_uva()

```

评论区精华

- reviewer ZJY0516 请求提供详细的 profile（如 nsys、torch profiler）以理解性能提升来源。作者贴出 profiler 截图显示该操作耗时从 ~0.8ms 降至 ~0.2ms。
- reviewer njhill 表示无法复现同等幅度的提升，认为属于微优化，并在 commit 中重构了代码（使用同一内核而非新增单独的内核）。作者回应称高并发（concurrency 40）下仍有明显收益。最终 njhill 批准合并。
- 性能提升幅度与可复现性 (performance): reviewer njhill 批准 PR，但性能收益未达成一致，作者的高并发场景数据被接受。

风险与影响

- 风险：
 1. 性能收益不确定性：reviewer 无法复现作者声称的显著提升，实际收益可能因硬件、模型、并发度而异。
 2. 多分支路径复杂性：新增 FusedStagedWriter 类及其与 apply_write 的切换逻辑，增加了代码维护成本；MULTI_GROUP 条件分支在 Triton 内核中可能因编译常量展开，但仍需测试验证两种路径的正确性。
 3. 测试覆盖局限：新增测试仅覆盖基本场景（2-3 个 group、2 个 request），缺少高压力、大并发、混合 group 数或伴随 num_blocks 同步变化的随机测试。
 4. 仅支持 CUDA：测试标记为 requires CUDA，其他平台（如 ROCm、CPU）可能无法运行或需要调整。
 - 影响：影响范围仅限于使用 V2 Model Runner 且具有多 KV-group（如 GQA/MQA）的模型。单 KV-group 模型无变化。收益为减少一次内核启动开销，理论上降低 execute_model 延迟，但实测效果因场景而异。对系统其他模块无影响，因为 BlockTables 属于内部数据结构，接口不变。团队需关注在目标模型上是否真有收益，并确保 ROCm/CPU 等后端的兼容性。
 - 风险标记：性能收益不确定，多路径分支复杂性，仅 CUDA 后端验证，测试覆盖较浅

关联脉络

- PR #42656 Apply LRU policy only to proper cache entries: 该 PR 优化了 KV cache 释放策略中的 staged writes 逻辑，与本次 block table 写入性能优化属于同一功能域（v1 block table / staged writes），但关注点不同。