

PR #50874 完整报告

vllm-project/vllm

[Bugfix][R3] Size monolithic routing replay buffer for DP

合并时间: 2026-08-13 15:18

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/50874>

执行摘要

- 一句话: 修复 DP/EP 下 MoE 路由回放缓冲容量与分片捕获
- 推荐动作: 值得精读: routing-replay 的布局契约 (DP/EP/SP 组合) 是典型易错点, `capture()` 的四种分支处理可作为并行布局分发的参考实现。重点关注 `local_sizes` 的 `all-gatherv` 布局映射、分支优先级、以及 `padding` 对缓冲预算的影响。

功能与动机

PR body 明确说明两个假设在 TP2/DP2 部署下失效: 其一, Replay buffer capacity——`max_num_tokens` 是 per-rank scheduler limit, 而 naive dispatch 路径在调用单体 kernel 前会 `all-gather` rank-local batches, 导致 8192 行缓冲面对 16384 行 kernel 输入并在 warmup 阶段失败; 其二, Gathered shard layout——DP+EP 下 `topk_ids` 可能包含带 CUDA-graph/SP padding 的 gathered sequence-parallel shards, 与按未 padding per-DP token 计数匹配的既有捕获路径冲突, 触发 batch-dimension assertion。

实现拆解

1. 缓冲分配扩容 (`vllm/model_executor/layers/fused_moe/modular_kernel.py`, `ModularExpert.set_capture_fn`): 新增 `dispatch_group_size = self.moe_config.ep_size if self.moe_config.use_ep else self.moe_config.dp_size`, 将 `_routing_replay_buffer` 第一维从 `max_num_tokens` 扩为 `max_num_tokens * dispatch_group_size`。原因是 naive dispatch 在调用单体 kernel 前会按 group 聚合各 rank batch, EP 开启时聚合的是扁平化 EP group 而非 DP group。
2. 新增第四种 batch 布局分支 (`vllm/model_executor/layers/fused_moe/routed_experts_capturer.py`, `RoutedExpertsCapturer.capture`): 读取 `ctx.dp_metadata.local_sizes` 并计算 `gathered_size`; 在既有 `n == total` 与 `n == token_num_per_dp` 分支之后插入 `shard_sizes is not None and n == gathered_size` 分支, 按 `all-gatherv` 的 DP-then-TP 扁平顺序计算本 DP rank 的连续 shard group 起止偏移, 仅复制 `token_num_per_dp` 行真实数据, 截掉 CUDA-graph/SP padding。
3. 文档与报错信息同步: `capture()` docstring 由三种布局扩展为四种, `AssertionError` 提示加入 `gathered_size`, 避免未来布局失配时难以定位。
4. 验证与配套: 未新增单元测试; 作者在 GB200 上以 TP2/DP2+EP 默认后端跑通 `mixed/decode` CUDA-graph capture 与 `/v1/completions` 请求, 并通过 pre-commit 全部 hooks。

关键文件：

- `vllm/model_executor/layers/fused_moe/routed_experts_capturer.py`（模块 路由捕获；类别 `source`；类型 `core-logic`；符号 `RoutedExpertsCapturer.capture`）：核心修复文件：`capture()` 新增 naive DP+EP 下基于 `dp_metadata.local_sizes` 的第四种 `batch` 布局分支，解决 `gathered shard` 含 `padding` 导致的 `batch` 维度断言失败。
- `vllm/model_executor/layers/fused_moe/modular_kernel.py`（模块 内核缓冲；类别 `source`；类型 `core-logic`；符号 `ModularExpert.set_capture_fn`）：`set_capture_fn` 中 `routing replay buffer` 从按单 `rank max_num_tokens` 分配改为按 `dispatch group`（`EP group` 或 `DP group`）大小分配，修复 `warmup` 阶段缓冲越界。

关键符号：`RoutedExpertsCapturer.capture`, `ModularExpert.set_capture_fn`

关键源码片段

`vllm/model_executor/layers/fused_moe/routed_experts_capturer.py`

核心修复文件：`capture()` 新增 naive DP+EP 下基于 `dp_metadata.local_sizes` 的第四种 `batch` 布局分支，解决 `gathered shard` 含 `padding` 导致的 `batch` 维度断言失败。

```
def capture(self, layer_id: int, topk_ids: torch.Tensor) -> None:
    # 根据 forward context 判断是否为多 DP 部署
    ctx = get_forward_context()
    if ctx.dp_metadata is None: # 单 DP：整张表就是本 rank 的 token
        start_loc = 0
        end_loc = topk_ids.shape[0]
        token_num_per_dp = topk_ids.shape[0]
    else:
        num_tokens_dp = ctx.dp_metadata.num_tokens_across_dp_cpu
        token_num_per_dp = int(num_tokens_dp[self.dp_rank].item())
        total = int(num_tokens_dp.sum().item())
        n = topk_ids.shape[0]
        # local_sizes 是 all-gatherv 的精确分片布局，包含 CUDA-graph / SP padding
        shard_sizes = getattr(ctx.dp_metadata, 'local_sizes', None)
        gathered_size = sum(shard_sizes) if shard_sizes is not None else -1

    if n == total:
        # naive dispatch: 所有 DP rank 的 token 拼接后再路由，
        # 本 rank 拥有区间 [end_loc - token_num_per_dp, end_loc)
        cumsum = torch.cumsum(num_tokens_dp, dim=0)
        end_loc = int(cumsum[self.dp_rank].item())
        start_loc = end_loc - token_num_per_dp
    elif n == token_num_per_dp:
        # modular-kernel 路径：DP combine 在 quant_method.apply 内部完成，
        # select_experts 只看得到本 rank 的 token
        start_loc = 0
        end_loc = token_num_per_dp
    elif shard_sizes is not None and n == gathered_size:
        # naive DP+EP dispatch: sequence-parallel shards 按 DP-then-TP
        # 顺序 flatten 到 EP group，local_sizes 是精确布局且含 padding.
```

```

# 定位本 DP rank 的连续 shard group, 再截掉尾部 padding。
num_dp_ranks = len(num_tokens_dp)
assert len(shard_sizes) % num_dp_ranks == 0
shards_per_dp_rank = len(shard_sizes) // num_dp_ranks
first_shard = self.dp_rank * shards_per_dp_rank
start_loc = sum(shard_sizes[:first_shard])
end_loc = start_loc + token_num_per_dp
elif (self.tp_size > 1
      and n != token_num_per_dp
      and n == (token_num_per_dp + self.tp_size - 1) // self.tp_size):
    # SP + modular-kernel 路径: 跨 TP group 沿 dim=0 all-gather,
    # 保留前 token_num_per_dp 行, 尾部是 SP ceil-div padding
    topk_ids = get_tp_group().all_gather(topk_ids, dim=0)
    start_loc = 0
    end_loc = token_num_per_dp
else:
    # 其余形状一律报错, 避免把数据写进错误的 buffer 区间
    sp_expected = ((token_num_per_dp + self.tp_size - 1) // self.tp_size
                   if self.tp_size > 0 else -1)
    raise AssertionError(
        'RoutedExpertsCapturer: unexpected topk_ids batch '
        f'dim {n} (expected {total}, {token_num_per_dp}, '
        f'{gathered_size}, or {sp_expected} for '
        f'dp_rank={self.dp_rank}, tp_size={self.tp_size})')

# 防御: 模型层数超过 buffer 维度时直接跳过
if layer_id >= self.device_buffer.shape[1]:
    return

# 只把本 DP rank 的真实 token 行写入 buffer, 截掉尾部 padding
self.device_buffer[:token_num_per_dp, layer_id, :] = topk_ids[
    start_loc:end_loc, :]

```

评论区精华

aoshen02 在 modular_kernel.py 的 set_capture_fn diff 上提了两条风格意见: 一是 "A bit verbose." (新增注释偏长); 二是建议把 `dispatch_group_size` 的三元表达式改为更清晰的 `if self.moe_config.use_ep: ... else: ...` 写法。第三个提交 "Address routing replay review feedback" 落实了这些反馈 (最终 head 版本注释精简为一行说明, EP 分支语义被采纳), 随后 ZJY0516 代 aoshen02 批准, ywang96 也 APPROVED。

- 缓冲分配注释过于冗长 (style): 第三个提交 "Address routing replay review feedback" 精简了注释, PR 随后获批。
- `dispatch_group_size` 取值方式可读性 (style): 最终 head 版本采纳了语义 (EP 开启时用 `ep_size`), 但保留了三元表达式写法; 提交 a4f624a 标注为处理 review 反馈, 随后获得 APPROVED。

风险与影响

- 风险：缓冲容量仍以 $\text{max_num_tokens} * \text{dispatch_group_size}$ 为预算，而 `local_sizes` 含 SP/CUDA-graph padding，极端重 padding 下 $\text{sum}(\text{local_sizes})$ 可能超出预算，届时 `_maybe_make_routing_replay_buffer` 的 `ValueError` 或 `buffer setitem` 会 fail-loud，需要在高 DP × 高 padding 场景复核。新分支依赖 `local_sizes` 的确切语义（all-gatherv 布局、DP-then-TP 顺序、shard 数可被 DP rank 数整除），语义变化会直接触发 `assert`。分支顺序敏感：`n == gathered_size` 判定位于 `n == total` 与 `n == token_num_per_dp` 之后，若 padding 使 `gathered_size` 恰好等于这两个值之一会走旧分支取错区间。此外未新增自动化测试，回归保护依赖 runtime validation。
- 影响：影响所有启用 routing-replay capture 的 FlashInfer 单体 MoE + naive DP/EP 部署（如 TP2/DP2+EP 的 GB200 配置），修复 warmup 崩溃与 batch 维度断言；single-DP、modular-kernel 路径行为不变。缓冲显存随 group size 线性增长（ $\text{int16} \times \text{experts_per_token} \times \text{group_size}$ ），量级很小（每 token 几字节）。对 MoE 路由回放功能与并行配置的兼容性是正向修复。
- 风险标记：核心路径变更，缺少自动化测试，缓冲容量对 padding 敏感，依赖 `local_sizes` 布局契约

关联脉络

- PR #44214 Add routing replay output for monolithic MoE (per PR body): PR body 指出该 PR 引入了 routing replay output，本 PR 是其 DP/EP 场景下的容量与布局修复。
- PR #48698 Scale FlashInfer B12x MoE workspace (per PR body): PR body 的 duplicate check 指出 #48698 只分离扩展 FlashInfer B12x MoE workspace，未触及本 PR 修复的 routing replay 输出分配。
- PR #50940 Edit routed_experts_capturer.py shape configuration (per PR body): 同样编辑 routed_experts_capturer.py 但只改 shape 配置，未处理 capture() 的 batch 布局分支，与本 PR 互补。
- PR #50759 Skip MoE padding inside router implementations (per PR body): 在 router 实现内跳过 MoE padding，位于 gathered-batch 布局处理的上游，不影响本 PR 的 all-gatherv 布局修复。