

PR #40589 完整报告

vllm-project/vllm

[Bugfix] Fix dp mtp hang

合并时间: 2026-07-07 05:08

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/40589>

执行摘要

- 一句话: 修复 DP+MTP 模式下因 drafter 不同步导致的 hang
- 推荐动作: 值得精读。该 PR 展示了多 rank 环境下因条件分歧导致集合通信 hang 的典型场景及修复思路。review 中的设计权衡讨论 (dummy_run vs 直接运行 drafter) 具有启发意义, 体现了对 drafter 合约边界的审慎处理。

功能与动机

PR body 说明: 在 dp=4, tp=2, mtp=3 配置下, 部分请求的 input_len+output_len 接近 max_model_len, 导致某些 DP rank 的 proposed_draft_token_ids 因 $\text{max_seq_len}(4098) + \text{num_spec_token}(3) > \text{max_model_len}(4100)$ 而被跳过, 而其他 rank 仍执行 drafter.dummy_run(), 从而引发 hang。

实现拆解

1. 重构 use_gpu_toks 判断: 将原有的 use_gpu_toks 复合条件拆分为 drafter_runs_model_forward 和 use_gpu_toks 两个变量, 提高可读性。
drafter_runs_model_forward 独立于 disable_padded_drafter_batch, 便于后续分支复用。
2. 调整 EAGLE/DraftModel 分支的 else 逻辑: 原 elif 改为 else, 使 valid_sampled_token_count_event 处理与 dummy_run 在同一层级, 确保无论是否调用 prepare_next_token_ids_padded, 都会在 DP>1 时执行 dummy_run。
3. 新增 DP 下的 dummy_run 调用: 在 input_fits_in_drafter 为 False 且 DP>1 时, 调用 self.drafter.dummy_run(num_tokens=1), 以匹配其他 rank 的集合通信次数, 防止 hang。
4. 重命名与清理: 将 propose_drafts_after_bookkeeping 重命名为 draft_after_bookkeeping, 并在最后分支中将赋值从 input_fits_in_drafter 改为固定 True (因为此时 input_fits_in_drafter 为 False, 原逻辑意图是只有 input_fits_in_drafter 时才能在 bookkeeping 后执行 drafter, 但实际逻辑有偏差, 修正为更清晰的语义)。
5. 在 draft_after_bookkeeping 分支中补充 dummy_run: 当 input_fits_in_drafter 为 True 但 bookkeeping 后需运行 drafter 时, 如果 DP>1, 也调用 dummy_run 以保持同步 (但实际条件下该分支 input_fits_in_drafter 为 True 时, 原 should_skip 不会触发, 因此 hang 不会发生, 这里修改更多是防御性)。

关键文件:

- vllm/v1/worker/gpu_model_runner.py (模块 模型运行器; 类别 source; 类型 core-logic)
: 核心修改文件, 修复 DP+MTP hang 的所有逻辑变更集中于此

关键符号: 未识别

关键源码片段

vllm/v1/worker/gpu_model_runner.py

核心修改文件, 修复 DP+MTP hang 的所有逻辑变更集中于此

```
# vllm/v1/worker/gpu_model_runner.py
```

```
def _execute_model(...):
    # ...
    spec_config = self.speculative_config
    draft_after_bookkeeping = False
    if spec_config is not None:
        input_fits_in_drafter = self._input_fits_in_drafter(
            spec_decode_common_attn_metadata)
        # 判断 drafter 是否执行 GPU 模型前向 (包含集合通信)
        drafter_runs_model_forward = (
            spec_config.use_eagle()
            or spec_config.uses_draft_model()
            or spec_config.uses_extract_hidden_states()
        )
        use_gpu_toks = (
            drafter_runs_model_forward
            and not spec_config.disable_padded_drafter_batch
        )
        if use_gpu_toks:
            # EAGLE/DraftModel 使用 GPU 采样 token 作为输入
            sampled_token_ids = sampler_output.sampled_token_ids
            if input_fits_in_drafter:
                propose_draft_token_ids(sampled_token_ids)
            else:
                # 超出 drafter 容量时, 零化 draft token
                if self.valid_sampled_token_count_event is not None:
                    # 准备填充后的采样 token ID
                    next_token_ids, valid_sampled_tokens_count = (
                        self.drafter.prepare_next_token_ids_padded( # 假设有该方法
                            sampled_token_ids, self.requests,
                            self.input_batch, self.discard_request_mask.gpu,
                        )
                    )
                self._copy_valid_sampled_token_count(
                    next_token_ids, valid_sampled_tokens_count)
        if self.parallel_config.data_parallel_size > 1:
            # 核心修复: 当其他 DP rank 执行了 dummy_run 时,
            # 本 rank 也需要执行一次 dummy_run 以匹配集合通信次数, 防止 hang
```

```

        self.drafter.dummy_run(num_tokens=1)
# ... 其他 speculative 方法分支 (ngram 等)
if not input_fits_in_drafter:
    # 零化 draft token 防止使用过期 draft
    self._draft_token_ids = torch.zeros(1, device=self.device, dtype=torch.long)
# ... 后续 bookkeeping 后 drafter 分支
if draft_after_bookkeeping:
    if input_fits_in_drafter:
        # 此时 input_fits_in_drafter 为 True, 正常执行 drafter
        propose_draft_token_ids(valid_sampled_token_ids)
    elif self.parallel_config.data_parallel_size > 1:
        # 防御性补充 dummy_run
        self.drafter.dummy_run(num_tokens=1)

```

评论区精华

Review 中 main reviewer MatthewBonanni 最初认为 dummy_run 不是正确做法, 因为其脆弱且不如直接运行 drafter (即使产生垃圾结果)。但最终他改变看法, 指出部分 drafter (如 rope_scaling) 在超出 max_model_len 时会直接崩溃, 不能假设它们能安全处理越界输入, 因此 dummy_run 是合理的解决方案。

- dummy_run vs 直接运行 drafter 的设计权衡 (design): 采用 dummy_run 方案, 因为不能假设所有 drafter 能安全处理越界输入。
- 变量作用域和类型安全问题 (correctness): 作者提交修复版本的 commit 解决了这些问题。

风险与影响

- 风险:
 1. 仅影响 DP>1 且使用 EAGLE/DraftModel 类 drafter 时触发的 else 分支, 不会影响非 DP 或非 GPU toks 的 speculative 方法 (如 ngram、draft model)。
 2. dummy_run(num_tokens=1) 开销小, 但需保证所有 drafter 类型均实现 dummy_run 方法。当前改动只对 use_gpu_toks 对应的 drafter 类型 (EagleProposer, DFlashProposer 等) 生效, 已确认这些 proposer 有 dummy_run 实现。
 3. 变量 input_fits_in_drafter 仅在 spec_config 不为 None 时定义, 但改动中所有引用都在该分支内, 因此不会出现 UnboundLocalError。- 影响: 直接影响: 使用 DP+MTP 且请求长度接近 max_model_len 的用户不再遇到 hang。影响范围: 仅限于启用 data_parallel_size>1 且使用 EAGLE/DraftModel 类 speculative decoding 的场景。性能影响极小, 因为 dummy_run 几乎无计算成本。- 风险标记: 核心路径变更, 仅影响 DP+MTP 场景

关联脉络

- PR #29422 相关 PR (MatthewBonanni 提及的之前尝试): MatthewBonanni 提到此前有类似意图的 PR 但被搁置, 虽然未合并, 但与本 PR 目标一致。