

PR #46757 完整报告

vllm-project/vllm

Fix Quark mxfp4 quantized model loading issue under mtp

合并时间: 2026-07-17 02:58

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/46757>

执行摘要

- 一句话: 修复 Quark MXFP4 模型 MTP 加载时权重形状不匹配
- 推荐动作: 建议合并, 修复了明确的加载错误。设计简洁, 逻辑集中在 `should_ignore_layer` 中。后续建议增加单元测试覆盖 `check_children` 分支, 并考虑在 `exclude` 配置中支持正则以更通用地解决此类问题。

功能与动机

GLM-5.2 MXFP4 检查点可将 MTP 层保持未量化, 而量化主 MoE 层。Quark 配置目前正确排除了线性层子层, 但 MTP draft 层构建为 FusedMoE 模块, `exclude` 未正确生效, 导致 vLLM 为未量化的 bf16 MTP 层分配 MXFP4-packed 专家权重, 加载时出现 `RuntimeError: The size of tensor a (256) must match the size of tensor b (512)`。

实现拆解

1. 在 `should_ignore_layer` 中新增 `check_children` 关键字参数, 当启用时遍历 `ignore` 列表, 若存在以 `layer_name` 开头 (或完全匹配) 的目标且非正则模式, 则判定该层应被忽略。
2. 在 `get_quant_method` 中为 `RoutedExperts` 层调用 `should_ignore_layer` 时传入 `check_children=True`; 若返回 `True`, 直接返回 `UnquantizedFusedMoEMethod` 而非继续走量化分支。
3. 同步更新 `quark.py` 的导入语句, 新增 `UnquantizedFusedMoEMethod`。
4. 验证方式: 通过 `vllm serve` 命令在 MI355X 上加载模型, 确认健康检查和 `completion` 正常, 且 `old MoE load failure` 不再出现。

关键文件:

- `vllm/model_executor/layers/quantization/quark/quark.py` (模块 量化层; 类别 `source`; 类型 `data-contract`; 符号 `get_quant_method`, `UnquantizedFusedMoEMethod`): 核心修改: 在 `get_quant_method` 中增加对 `RoutedExperts` 被忽略时的处理, 并导入 `UnquantizedFusedMoEMethod`。
- `vllm/model_executor/layers/quantization/quark/utils.py` (模块 量化工具; 类别 `source`; 类型 `data-contract`; 符号 `should_ignore_layer`): 核心逻辑: 在 `should_ignore_layer` 中新增 `check_children` 参数和子层检查逻辑, 支持父层因子层被忽略而整体忽略。

关键符号: `should_ignore_layer`, `get_quant_method`

关键源码片段

vllm/model_executor/layers/quantization/quark/quark.py

核心修改：在 `get_quant_method` 中增加对 `RoutedExperts` 被忽略时的处理，并导入 `UnquantizedFusedMoEMethod`。

```
# SPDX-License-Identifier: Apache-2.0
from vllm.model_executor.layers.fused_moe import (
    RoutedExperts,
    UnquantizedFusedMoEMethod, # 新增导入，用于返回未量化的 MoE 方法
)

class QuarkConfig(...):
    def get_quant_method(
        self, layer: torch.nn.Module, prefix: str
    ) -> 'QuantizeMethodBase | None':
        # 检查该层是否应该被跳过量化
        exclude_layers = cast(list[str], self.quant_config.get('exclude'))
        if should_ignore_layer(
            prefix,
            ignore=exclude_layers,
            fused_mapping=self.packed_modules_mapping,
            check_children=isinstance(layer, RoutedExperts), # 对 RoutedExperts 启用子层检查
        ):
            # 如果是 RoutedExperts 且被忽略，直接返回未量化的 MoE 方法
            if isinstance(layer, RoutedExperts):
                return UnquantizedFusedMoEMethod(layer.moe_config)
            # 其他层的处理 (self_attn, LinearBase) 保持不变 ...
```

vllm/model_executor/layers/quantization/quark/utils.py

核心逻辑：在 `should_ignore_layer` 中新增 `check_children` 参数和子层检查逻辑，支持父层因子层被忽略而整体忽略。

```
def should_ignore_layer(
    layer_name: str | None,
    ignore: Iterable[str],
    fused_mapping: Mapping[str, list[str]] = MappingProxyType({}),
    *,
    check_children: bool = False, # 新增关键字参数，默认不检查子层
) -> bool:
    if layer_name is None:
        return False

    # MoE 层目前是全或无的：如果有任何子层被忽略，父层也必须被忽略
    if check_children and any(
        target == layer_name or target.startswith(layer_name + '.')
        for target in ignore
    ):
        if not target.startswith('re:') # 正则表达式模式不经过此检查
    ):
        return True
    return layer_name in ignore
```

```
return True
```

```
# 后续处理 fused_mapping 等（保持不变）
```

```
# ...
```

评论区精华

- BowenBao 提议将忽略逻辑移入 `should_ignore_layer`，作者采纳并重构了代码。
- BowenBao 后续建议简化变量 `is_ignored` 直接使用 `if should_ignore_layer(...)`，作者按建议修改。
- BowenBao 要求为子层检查逻辑添加注释，作者添加了详细注释说明 MoE 层 all-or-nothing 假设。
- fxmarty-amd 指出理想方案应在 `exclude` 列表中使用正则表达式，当前 `any` 检查假设所有子模块均在 `exclude` 列表中，不够安全，但作为快速修复可接受。
- 将 `RoutedExperts` 忽略逻辑移入 `should_ignore_layer (design)`: 作者同意并将逻辑移至 `should_ignore_layer`，新增 `check_children` 参数。
- 简化 `is_ignored` 变量使用 (style): 作者按建议修改，去除了中间变量。
- 为子层检查添加注释 (documentation): 作者在 `should_ignore_layer` 中添加了详细注释，解释 MoE 层 all-or-nothing 的假设及参考的 config 链接。
- 考虑使用正则表达式替代 `any` 检查 (design): 当前 PR 保留 `any` 检查，认为在现有配置下有效；未来可考虑支持正则。

风险与影响

- 风险：风险较低：仅影响 Quark 量化下 `RoutedExperts` 层的 `ignore` 判断逻辑。
`check_children` 的 `any` 检查假设 `exclude` 列表中同时包含所有相关子层，若部分子层不在列表中可能导致错误忽略；但实际配置中通常全部列出或使用正则。未新增单元测试，手动测试覆盖有限，存在回归可能。
- 影响：影响范围小：仅影响使用 Quark 量化且模型包含 `RoutedExperts` (MoE) 层并配置了子层 `exclude` 的用户，典型场景是 GLM-5.2 MTP 推理。对无量化或非 MoE 模型无影响。
修复后 MTP 加载正常，可通过 TP=4 验证。
- 风险标记：缺少测试覆盖，子层假设可能不完全

关联脉络

- 暂无明显关联 PR