

PR #47832 完整报告

vllm-project/vllm

[CI] Fix Transformers modeling backend LoRA test

合并时间: 2026-07-07 19:40

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/47832>

执行摘要

- 一句话: 修复 LoRA 下 QKV 融合因缺少 `split_sizes` 属性导致的崩溃
- 推荐动作: 值得精读, 设计决策 (用已有属性推导而非存储冗余) 值得借鉴。

功能与动机

PR #47804 和 #47824 指出的 QA 测试失败根本原因是 #47187 新增的 `split_sizes` 属性未被 `LoRA` 包装类复制。作者选择从根本上消除冗余属性, 降低维护成本。

实现拆解

1. 移除 `split_sizes` 属性: 在 `QKVParallelLinear.__init__` 中删除手动设置 `split_sizes` 的代码, 改为直接使用已有的 `output_sizes` 和 `tp_size` 属性。
2. 修改 `QKVFuser` 的 AST 模板: 在 `qkv.py` 中将 `split` 调用从 `self.qkv.split_sizes` 改为 `[s // self.qkv.tp_size for s in self.qkv.output_sizes]`, 避免依赖新属性。
3. 更新测试存根: 在 `test_linear.py` 中将存根对象的 `split_sizes` 替换为 `output_sizes` 和 `tp_size`, 并验证生成的代码包含 `tp_size` 引用。

关键文件:

- `vllm/model_executor/models/transformers/fusers/qkv.py` (模块 融合器; 类别 `source`; 类型 `core-logic`; 符号 `update_forward`, `update_attrs`): 核心 AST 模板变更, 用 `output_sizes/tp_size` 替换 `split_sizes`
- `vllm/model_executor/linear.py` (模块 线性层; 类别 `source`; 类型 `data-contract`; 符号 `QKVParallelLinear.init`): `QKVParallelLinear` 初始化简化为无重复 `output_size` 计算
- `tests/models/transformers/fusers/test_linear.py` (模块 线性层; 类别 `test`; 类型 `test-coverage`; 符号 `_apply_qkv_fuser_with_stubs`, `test_detects_and_rewrites_qkv`): 测试存根和断言适配新方案

关键符号: `QKVParallelLinear.init`, `QKVFuser.update_forward`, `QKVFuser.update_attrs`, `_apply_qkv_fuser_with_stubs`

关键源码片段

`vllm/model_executor/models/transformers/fusers/qkv.py`

核心 AST 模板变更, 用 `output_sizes/tp_size` 替换 `split_sizes`

```
# vllm/model_executor/models/transformers/fusers/qkv.py
# 在 update_forward 中，字符串模板从 split_sizes 改为动态计算：
# 之前：self.qkv.split_sizes，依赖手动维护的属性
# 之后：自 output_sizes 和 tp_size 推导，彻底消除冗余
sections = f"[s // {merged}.tp_size for s in {merged}.output_sizes]"
template = f"{'', '.join(temps)} = {merged}(__arg__).split({sections}, -1)"

# 同时删除了 update_attrs 中设置 split_sizes 的代码块
# 因为 QKVParallelLinear 的 output_sizes 已包含所有必要信息
```

vllm/model_executor/layers/linear.py

QKVParallelLinear 初始化简化为无重复 output_size 计算

```
# vllm/model_executor/layers/linear.py
class QKVParallelLinear(RowParallelLinear):
    def __init__(self, ...):
        # ... 初始化属性 ...
        # 之前：手动计算 output_size + 设置 output_sizes
        # output_size = (num_heads * head_size + ...) * tp_size
        # self.output_sizes = [...]

        # 之后：先设置 output_sizes，再通过 sum 得到 output_size
        self.output_sizes = [
            self.num_heads * self.head_size * tp_size, # q_proj
            self.num_kv_heads * self.head_size * tp_size, # k_proj
            self.num_kv_heads * self.v_head_size * tp_size, # v_proj
        ]
        output_size = sum(self.output_sizes)

        super().__init__(...)
```

tests/models/transformers/fusers/test_linear.py

测试存根和断言适配新方案

```
# tests/models/transformers/fusers/test_linear.py
# 之前：merged.split_sizes = [q.out_features, k.out_features, v.out_features]
# 之后：使用 output_sizes 和 tp_size 模拟真实类行为
merged.output_sizes = [q.out_features, k.out_features, v.out_features]
merged.tp_size = 1

# 测试断言更新：验证编译后的代码包含 output_sizes 和 tp_size
names = code.co_names
assert "qkv_proj" in names and "output_sizes" in names and "o_proj" in names
assert "tp_size" in names
```

评论区精华

无 review 讨论，仅 [claude\[bot\]](#) 自动评论说明分支来自 fork 未启用自动 review。

- 暂无高价值评论线程

风险与影响

- 风险：移除属性可能影响其他依赖 `split_sizes` 的代码（如需提前访问该属性的序列化或工具函数），但当前库内无其他使用者。性能方面，动态计算 `tp_size` 除法的开销在 `eager` 模式极小，且会被 `torch.compile` 消除。
- 影响：直接修复 LoRA 测试回归，使 Transformers 后端融合功能正常可用。影响范围限于 `QKVParallelLinear` 和 `QKVFuser` 使用方，无用户可见变更。
- 风险标记：属性移除可能影响外部扩展

关联脉络

- PR #47804 [BugFix] LoRA test fix for qkv merge: 此 PR 的原尝试，被本 PR 取代
- PR #47824 [Bugfix] Add qkv parallel linear lora prefix for lora: 另一类似修复尝试，被本 PR 取代
- PR #47187 QKV fusion for Transformers backend: 引入 `split_sizes` 属性导致回归的原始 PR