

PR #43981 完整报告

vllm-project/vllm

[AMD][Bugfix][Quantization] Honor fused-name match in is_layer_skipped

合并时间: 2026-06-16 00:37

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/43981>

执行摘要

- 一句话: 修复 is_layer_skipped 未处理 fused name 直接匹配的回归
- 推荐动作: 值得精读, 特别是参与量化或模型配置的开发者。本 PR 清晰展示了如何处理量检查点格式兼容性问题, 设计决策明确: 优先检查 fused name 自身匹配, 再 fallback 到 per-shard。测试代码详尽, 可作为类似场景的参考。

功能与动机

PR #41892 在修复 Quark INT8 检查点的同时, 引入回归: 对于直接列出 fused name 的 FP8 检查点, is_layer_skipped 错误地展开 shard, 导致本应跳过的层被量化, 输出退化。本 PR 修复此逻辑, 优先检查 fused name 匹配, 兼容两种检查点格式。

实现拆解

1. 定位问题: vllm/model_executor/layers/quantization/utils/quant_utils.py 中的 is_layer_skipped 函数, 当 proj_name 在 fused_mapping 中时, 总是将 fused name 展开为多个 unfused shard 前缀, 忽略了 fused name 本身可能已在 ignored_layers 中的情况。
2. 修改匹配顺序: 新增短路检查: 若 proj_name 在 fused_mapping 中且 match_func(prefix, ignored_layers) 成立, 则直接返回 True。只有不匹配时才进入 shard 展开逻辑。这保留了原有 Quark INT8 路径的行为。
3. 添加单元测试: 在 tests/quantization/test_quark.py 中新增 6 个测试函数, 覆盖 fused name 直接匹配、unfused shard 全匹配、部分 shard 引发异常、无匹配、非 fused 层不变以及子串匹配模式。确保新逻辑正确且无回归。

关键文件:

- vllm/model_executor/layers/quantization/utils/quant_utils.py (模块 量化工具; 类别 source; 类型 core-logic; 符号 is_layer_skipped): 核心逻辑变更, 修复 fused name 直接匹配问题
- tests/quantization/test_quark.py (模块 量化测试; 类别 test; 类型 test-coverage; 符号 test_fused_name_listed_directly_is_skipped, test_unfused_shards_listed_is_skipped, test_partial_shards_raises, test_not_skipped_when_nothing_listed): 新增 6 个单元测试覆盖 fused name 匹配逻辑

关键符号: is_layer_skipped

关键源码片段

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

核心逻辑变更，修复 fused name 直接匹配问题

```
def is_layer_skipped(
    prefix: str,
    ignored_layers: list[str],
    fused_mapping: Mapping[str, list[str]] = MappingProxyType({}),
    *,
    skip_with_substr: bool = False,
) -> bool:
    def prefix_full_match(prefix: str, ignored_layers: list[str]) -> bool:
        return prefix in ignored_layers

    def substr_match(prefix: str, ignored_layers: list[str]) -> bool:
        return any(layer in prefix for layer in ignored_layers)

    match_func = substr_match if skip_with_substr else prefix_full_match

    proj_name = prefix.split(".")[1]

    # 新逻辑：如果 proj_name 在 fused_mapping 中且 prefix 本身匹配忽略列表，则直接跳过
    if proj_name in fused_mapping and match_func(prefix, ignored_layers):
        is_skipped = True
    elif proj_name in fused_mapping:
        # 只有 fused name 不匹配时才展开 shard
        shard_prefixes = [
            prefix.replace(proj_name, shard_proj_name)
            for shard_proj_name in fused_mapping[proj_name]
        ]

        is_skipped = None
        for shard_prefix in shard_prefixes:
            is_shard_skipped = match_func(shard_prefix, ignored_layers)
            if is_skipped is None:
                is_skipped = is_shard_skipped
            elif is_shard_skipped != is_skipped:
                raise ValueError(
                    f"Detected some but not all shards of {prefix} "
                    "are quantized. All shards of fused layers "
                    "to have the same precision."
                )
    elif "experts" in prefix and not skip_with_substr:
        expert_ignore_layers = filter(
            lambda layer_name: "experts" in layer_name, ignored_layers
        )
    return any(
        prefix in layer_name if not skip_with_substr else layer_name in prefix
    )
```

```

        for layer_name in expert_ignore_layers
    )
else:
    is_skipped = match_func(prefix, ignored_layers)

assert is_skipped is not None
return is_skipped

```

tests/quantization/test_quark.py

新增 6 个单元测试覆盖 fused name 匹配逻辑

```

FUSED_MAPPING = {
    "qkv_proj": ["q_proj", "k_proj", "v_proj"],
    "gate_up_proj": ["gate_proj", "up_proj"],
}

def test_fused_name_listed_directly_is_skipped():
    # 回归测试: FP8 检查点直接列出 fused name
    ignored = ["model.layers.0.self_attn.qkv_proj"]
    assert is_layer_skipped(
        prefix="model.layers.0.self_attn.qkv_proj",
        ignored_layers=ignored,
        fused_mapping=FUSED_MAPPING,
    )
    assert is_layer_skipped(
        prefix="model.layers.0.mlp.gate_up_proj",
        ignored_layers=["model.layers.0.mlp.gate_up_proj"],
        fused_mapping=FUSED_MAPPING,
    )

def test_unfused_shards_listed_is_skipped():
    # Quark INT8 风格: 列出所有 unfused shard
    ignored = [
        "model.layers.0.self_attn.q_proj",
        "model.layers.0.self_attn.k_proj",
        "model.layers.0.self_attn.v_proj",
    ]
    assert is_layer_skipped(
        prefix="model.layers.0.self_attn.qkv_proj",
        ignored_layers=ignored,
        fused_mapping=FUSED_MAPPING,
    )

def test_partial_shards_raises():
    ignored = ["model.layers.0.self_attn.q_proj"]
    with pytest.raises(ValueError):

```

```

is_layer_skipped(
    prefix="model.layers.0.self_attn.qkv_proj",
    ignored_layers=ignored,
    fused_mapping=FUSED_MAPPING,
)

def test_not_skipped_when_nothing_listed():
    assert not is_layer_skipped(
        prefix="model.layers.0.self_attn.qkv_proj",
        ignored_layers=["model.layers.0.mlp.gate_up_proj"],
        fused_mapping=FUSED_MAPPING,
    )

def test_non_fused_layer_unaffected():
    assert is_layer_skipped(
        prefix="model.layers.0.self_attn.o_proj",
        ignored_layers=["model.layers.0.self_attn.o_proj"],
        fused_mapping=FUSED_MAPPING,
    )
    assert not is_layer_skipped(
        prefix="model.layers.0.self_attn.o_proj",
        ignored_layers=["model.layers.1.self_attn.o_proj"],
        fused_mapping=FUSED_MAPPING,
    )

def test_substr_match_on_fused_name():
    assert is_layer_skipped(
        prefix="model.layers.0.self_attn.qkv_proj",
        ignored_layers=["self_attn.qkv_proj"],
        fused_mapping=FUSED_MAPPING,
        skip_with_substr=True,
    )

```

评论区精华

- 要求单元测试: Reviewer okorzh-amd 要求为新增逻辑添加单元测试。作者询问后, okorzh-amd 明确为“逻辑的单元测试”。作者在 `test_is_layer_skipped.py` 中添加, 后合并到 `test_quark.py`。
- 测试文件位置: [tjtanaa](#) 询问 [hmellor](#) 测试文件的最佳位置, [hmellor](#) 建议归入已有的 `test_quark.py`, 作者完成移动。
- 要求添加单元测试 (testing): 测试已添加并整合到 `test_quark.py` 中, okorzh-amd 认可 (“Looks good. thanks”)。
- 测试文件放置位置 (testing): 测试从独立的 `test_is_layer_skipped.py` 移动到 `test_quark.py`, 与其他量化测试一起。

风险与影响

- 风险：风险较低：变更高度局部化，仅在 `is_layer_skipped` 中增加一个短路判断。新分支不影响已有 Quark INT8 路径。新增的 6 个单元测试覆盖了主要场景。但需注意未来若有新的 `fused_mapping` 定义或不同的忽略列表格式，可能需要进一步调整。
- 影响：影响范围：仅影响具有 `packed_modules_mapping` 且其检查点直接列出 `fused name` 的模型（如 Step-3.5-Flash-FP8）。修复了这些用户的推理错误。系统无性能影响。
团队：该修复提高了量化跳过逻辑的鲁棒性。
- 风险标记：回归风险，核心路径变更，测试覆盖改进

关联脉络

- PR #41892 [Bugfix][Quark] Fix W8A8 INT8 garbage outputs on Step-3.5-Flash (and other 3-key fused-MoE Quark exports): 本 PR 修复了 #41892 引入的回归。该 PR 添加了 `packed_modules_mapping` 但未考虑 `fused name` 直接匹配的情况。