

PR #49329 完整报告

vllm-project/vllm

[ROCm][CI] Fix order-dependent failure in test_flash_attn_accepts_handled_fp8_variants (MI355)

合并时间: 2026-07-22 07:54

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/49329>

执行摘要

- 一句话: 修复 ROCm 测试因模块绑定顺序导致的条件失败
- 推荐动作: 值得合并: 这是一个精准的测试 bug 修复, 有清晰的根因分析和最小改动方案。可作为测试中模块导入和 mock 作用的典型案例阅读。

功能与动机

在 `Kernels (B200-MI355)` CI 组中全量运行 `test_attention_selector.py` 时, `test_flash_attn_accepts_handled_fp8_variants[fp8!fp8_e4m3]` 用例失败, 但单独运行通过。失败原因是测试对 `is_xpu` 的 patch 目标模块错误, 导致在特定测试顺序下无法生效。该修复确保 ROCm 上的测试环境能正确验证 `FlashAttentionBackend` 对 `fp8 dtype` 的接受行为, 消除 CI 误报。

实现拆解

1. 定位根因: 在 `vllm/v1/attention/backends/fa_utils.py` 中, `flash_attn_supports_kv_cache_dtype` 函数通过 `current_platform.is_xpu()` 判断 XPU 平台, `current_platform` 是 `fa_utils` 模块级别的绑定。原测试通过 `import vllm.v1.attention.backends.flash_attn as fa_mod` 导入并 patch 了 `fa_mod.current_platform.is_xpu`, 但 `FlashAttentionBackend.supports_kv_cache_dtype` 实际调用的是 `fa_utils` 模块中的函数, 因此 `fa_mod` 的 patch 无影响。
2. 修改 patch 目标: 将导入改为 `import vllm.v1.attention.backends.fa_utils as fa_utils_mod`, 并 patch `fa_utils_mod.current_platform.is_xpu`, 确保 patch 作用于实际读取的模块绑定。
3. 添加注释: 在代码中添加了详细注释解释 patch 目标和原因, 提高可维护性。
4. 验证: 在 MI355 上全量运行 `tests/kernels/attention/test_attention_selector.py`, 全部 23 个测试通过, 7 个跳过; 两参数化用例也单独通过。

关键文件:

- `tests/kernels/attention/test_attention_selector.py` (模块 注意力测试; 类别 test; 类型 test-coverage): 唯一变更文件, 修复测试中的模块导入和 patch 目标, 消除顺序依赖失败。

关键符号: `test_flash_attn_accepts_handled_fp8_variants`

关键源码片段

tests/kernels/attention/test_attention_selector.py

唯一变更文件，修复测试中的模块导入和 patch 目标，消除顺序依赖失败。

```
# tests/kernels/attention/test_attention_selector.py
@pytest.mark.parametrize("kv_cache_dtype", ["fp8", "fp8_e4m3"])
def test_flash_attn_accepts_handled_fp8_variants(
    kv_cache_dtype: str, monkeypatch: pytest.MonkeyPatch
):
    """FlashAttentionBackend must accept the two fp8 dtypes it can actually
    handle: 'fp8' (alias for fp8_e4m3fn) and 'fp8_e4m3'."""
    # 关键修复: fp8 接受判断在 fa_utils 中使用其自身的 current_platform 绑定
    # 因此需要 patch fa_utils 的绑定, 而非 flash_attn 的绑定,
    # 以保持对前面测试可能替换 vllm.platforms.current_platform 的健壮性
    import vllm.v1.attention.backends.fa_utils as fa_utils_mod
    from vllm.v1.attention.backends.flash_attn import FlashAttentionBackend
    monkeypatch.setattr(fa_utils_mod.current_platform, "is_xpu", lambda: True)
    assert FlashAttentionBackend.supports_kv_cache_dtype(kv_cache_dtype)
```

评论区精华

无，仅一个批准评论 **LGTM**。无争议。

- 暂无高价值评论线程

风险与影响

- 风险：风险极低：本次修改仅涉及测试代码，且 patch 目标从错误的模块绑定改为正确的模块绑定，测试后 monkeypatch 自动还原，不会影响其他测试或生产环境。唯一细微风险是未来 fa_utils 模块内 current_platform 的引用方式变更可能导致失效，但概率很低。
- 影响：仅影响 ROCm 平台（MI355）上的 CI 测试执行。修复了顺序依赖导致的假阴性失败，消除 CI 随机失败。无用户或系统功能影响。
- 风险标记：暂无

关联脉络

- PR #46080 [Hardware][AMD][CI] Fix Kernels Attention test groups: 该 PR 将 test_attention_selector.py 完整引入 Kernels (B200-MI355) CI 组，并导致新的测试顺序，暴露了本 PR 修复的模块绑定问题。
- PR #42685 Original FP8 test addition (implied): 原测试用例（patch flash_attn 模块）由该 PR 引入，本 PR 修正了其 patch 目标。