

PR #46298 完整报告

vllm-project/vllm

[Hardware][AMD][CI] Fix gfx942 Kernels MoE test group

合并时间: 2026-06-22 05:45

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/46298>

执行摘要

- 一句话: 修复 AMD gfx942 的 MoE 测试组
- 推荐动作: 建议快速合并, 这是典型的 CI 修复工作, 逻辑清晰。可关注后续 gfx950 的完整修复 PR, 以全面覆盖 AMD MoE 测试。

功能与动机

PR body 明确指出: 'This PR fixes and gates the gfx942-based Kernels MoE test group. A full fix for this test group on gfx950 is still in progress.' 目的 (Purpose) 是修复和限定 gfx942 的 MoE 测试组。

实现拆解

1. 修复测试文件(`tests/kernels/moe/test_ocp_mx_moe.py`): 将 `Llama-4-Scout-17B-16E-Instruct-2-layers-mx_fp4` 模型从 `fxmarty` 仓库改为 `mawong-amd` 仓库 (因 `config.json` 验证错误, 作者自述模型验证修复); 导入方式从 `vllm._aiter_ops.rocm_aiter_ops` 改为 `vllm._aiter_ops.is_aiter_found`, 消除对内部模块的依赖; 函数调用 `convert_to_mxfp4_moe_kernel_format` 和 `shared_experts=None` 分别更新为 `convert_gpt_oss_weight_to_mxfp4_moe_kernel_format` 与 `layer=None`, 以及添加 `gpu_memory_utilization=0.8` 以应对 mxfp6 模型的大显存占用。
2. 配置 AMD 专用 CI 任务(`.buildkite/test_areas/kernels.yaml`): 为 `Kernels MoE Test %N` 任务添加 `mirror.amd` 子配置, 指定设备为 `mi325_1` (对应 gfx942), 超时 50 分钟, 额外依赖 `vllm/_aiter_ops.py` 和 `vllm/platforms/rocm.py`。此配置仅对 AMD 镜像生效, 不影响其他厂商。
3. 优化 AMD 测试参数(`.buildkite/test-amd.yaml`): 将 gfx942 和 gfx950 两个 MoE 任务的超时从 180 分钟降至 50 分钟 (实测足够), 并行度从 4 增至 5, 并将 gfx942 任务标记为 `optional: true` (允许非阻塞失败)。

关键文件:

- `tests/kernels/moe/test_ocp_mx_moe.py` (模块 MoE 测试; 类别 test; 类型 test-coverage): 测试逻辑调整核心: 修复模型 ID、简化导入、更新函数调用和增加内存参数, 确保测试在 gfx942 上通过。
- `.buildkite/test_areas/kernels.yaml` (模块 CI 配置; 类别 config; 类型 configuration): 为 AMD 添加 mirror 配置, 指定 gfx942 设备和额外依赖, 确保 CI 仅在该硬件上运行并能

正确检测触发条件。

- .buildkite/test-amd.yaml (模块 CI 配置; 类别 config; 类型 configuration) : 调整 gfx942 和 gfx950 MoE 测试的超时、并行度和 optional 标记, 优化 CI 资源利用。

关键符号: 未识别

关键源码片段

tests/kernels/moe/test_ocp_mx_moe.py

测试逻辑调整核心: 修复模型 ID、简化导入、更新函数调用和增加内存参数, 确保测试在 gfx942 上通过。

```
# tests/kernels/moe/test_ocp_mx_moe.py
# 导入简化: 从 vllm._aiter_ops 直接导入 is_aiter_found,
# 不再依赖 rocm_aiter_ops 内部模块, 使导入路径更直接。
from vllm._aiter_ops import is_aiter_found

# 模块级标志: 使用统一的 is_aiter_found() 来判断 aiter 可用性,
# 替代原来的 from vllm._aiter_ops import rocm_aiter_ops; ROCM_AITER_AVAILABLE = rocm_
aiter_ops.is_enabled()
ROCM_AITER_AVAILABLE = is_aiter_found()

# 测试用例参数化列表: 模型 ID 从 fxmarty 改为 mawong-amd,
# 因为原始模型 config.json 验证失败, 新模型仅修复了配置文件, 权重不变。
@pytest.mark.parametrize(
    "model_case",
    [
        ModelCase("fxmarty/qwen_1.5-moe-a2.7b-mx4", tp=2),
        ModelCase("fxmarty/deepseek_r1_3_layers_mx4", tp=8),
        ModelCase("mawong-amd/Llama-4-Scout-17B-16E-Instruct-2-layers-mx4", tp=1),
        ModelCase("fxmarty/Llama-3.1-70B-Instruct-2-layers-mx6", tp=1),
        ModelCase("fxmarty/Llama-3.1-70B-Instruct-2-layers-mx6", tp=4),
    ],
)

# 测试函数: 额外增加 gpu_memory_utilization=0.8 以避免 mx6 模型显存溢出
def test_mx4_loading_and_execution_moe(vllm_runner, model_case: ModelCase):
    ...
    with vllm_runner(
        model_case.model_id,
        tensor_parallel_size=model_case.tp,
        load_format="dummy",
        compilation_config={"cudagraph_capture_sizes": [16]},
        gpu_memory_utilization=0.8, # mx6 models use more scratch space
    ) as llm:
        ...
```

评论区精华

仅有一条 reviewer 评论：mawong-amd 在模型 ID 变更处解释 'Fixed a model validation error (`config.json`), everything else in this model remains the same otherwise', 即模型权重和逻辑不变，仅修复了配置文件验证问题。无其他讨论。

- 模型 ID 变更原因 (question): 模型 ID 变更仅修复配置问题，不改变模型权重或行为。

风险与影响

- 风险：
 - 回归风险（低）：测试模型 ID 变更仅影响 CI 测试，不涉及运行时逻辑；导入和函数调用变更均保持语义等价。
 - CI 覆盖风险（中）：gfx950 的 MoE 测试问题仍未解决，测试组被标记为 optional，可能降低 gfx950 的问题可见性。
 - 扩展性风险（低）：硬编码设备 mi325_1 可能随硬件代际更新而需要调整，但当前符合目标。
- 影响：
 - 用户影响：无直接影响，仅 CI 相关。
 - 系统影响：AMD gfx942 CI 中的 MoE 测试组将回归绿色，提高 CI 稳定性。
 - 团队影响：为 AMD 平台的 MoE 测试提供可靠的 CI 门禁，方便后续迭代。
 - 风险标记：gfx950 未覆盖，CI 测试可跳过

关联脉络

- PR #46080 [Hardware][AMD][CI] Fix Kernels Attention test groups: 同为 AMD CI 修复，针对不同测试组（Attention vs MoE），模式类似。
- PR #45967 [ROCm][CI] skip test_double_aiter_rms_quant_fusion: 同为 ROCm CI 修复，涉及测试跳过和配置调整。