

PR #46859 完整报告

vllm-project/vllm

[Hardware][AMD][CI] Fix Kernels Quantization test timeout

合并时间: 2026-06-27 04:19

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/46859>

执行摘要

- 一句话: 修复 AMD CI NVFP4 量化测试超时, 临时移除不稳定模型
- 推荐动作: 该 PR 值得合并。建议后续跟踪 `nvidia/Kimi-K2.6-NVFP4` 下载稳定性问题, 待网络问题解决后重新添加该模型配置。fixture 设计值得在其他 CI 测试中推广。

功能与动机

AMD CI 的 Kernels Quantization 测试组因 `test_nvfp4_moe_correctness` 下载 `nvidia/Kimi-K2.6-NVFP4` 模型的 shard 文件超时, 导致整个测试套件挂起 6 小时。PR #44667 合并后引入该问题, `nvidia/Kimi-K2.6-NVFP4` 的单个文件 (如 `model-00002-of-00060.safetensors`) 下载不稳定。

实现拆解

1. 添加 `MOE_MODEL_CONFIGS` 字典: 在文件顶部集中定义模型配置 (shard 文件列表、专家前缀等), 替代原来分散在多个函数中的硬编码。
2. 创建 `loaded_model_files` fixture: 使用 pytest 的 `scope="module"` fixture, 在模块级别执行 `huggingface_hub.snapshot_download`, 确保整个模块内的测试共享相同的下载路径, 避免重复下载。
3. 更新测试函数: `test_triton_dequantize_nvfp4` 和 `test_nvfp4_moe_correctness` 改为接收 `loaded_model_files` fixture, 并使用 `MOE_MODEL_CONFIGS` 获取 shard 路径和专家前缀。
4. 移除 `Kimi-K2.6-NVFP4` 模型: 从 `MOE_MODEL_CONFIGS` 字典中删除该模型条目, 暂时规避其下载超时问题, 待网络问题解决后可重新添加。
5. 清理重复下载逻辑: `_load_nvfp4_moe_weights` 函数也改为使用 fixture 提供的路径, 消除直接调用 `huggingface_hub.snapshot_download` 的冗余代码。

关键文件:

- `tests/kernels/quantization/test_nvfp4_emulation.py` (模块 量化测试; 类别 test; 类型 test-coverage; 符号 `loaded_model_files`, `test_triton_dequantize_nvfp4`): 唯一变更文件, 重构模型下载逻辑, 添加 fixture, 移除不稳定的模型配置

关键符号: `loaded_model_files`, `test_triton_dequantize_nvfp4`

关键源码片段

tests/kernels/quantization/test_nvfp4_emulation.py

唯一变更文件，重构模型下载逻辑，添加 fixture，移除不稳定的模型配置

集中管理模型下载配置，避免硬编码

```
MOE_MODEL_CONFIGS = {
    "nvidia/Qwen3-30B-A3B-NVFP4": {
        "shards": ["model-00001-of-00004.safetensors"],
        "expert_prefix": "model.layers.9.mlp.experts.",
        "expert_idx_pos": 5,
    }
}

@pytest.fixture(scope="module")
def loaded_model_files():
    # 模块级 fixture：只下载一次，所有测试共享
    return {
        model_id: huggingface_hub.snapshot_download(
            repo_id=model_id, allow_patterns=config["shards"]
        )
        for model_id, config in MOE_MODEL_CONFIGS.items()
    }
```

测试函数改为接受 fixture，避免重复下载

```
def test_triton_dequantize_nvfp4(monkeypatch, loaded_model_files) -> None:
    checkpoint_path = loaded_model_files["nvidia/Qwen3-30B-A3B-NVFP4"]
    shards = cast(list[str], MOE_MODEL_CONFIGS["nvidia/Qwen3-30B-A3B-NVFP4"]["shards"])
    shard_path = f"{checkpoint_path}/{shards[0]}"
    # ... 其余测试逻辑不变
```

评论区精华

该 PR 无 review 评论，仅由 Claude bot 自动评论（来自 fork，未触发审核），以及 maintainer AndreasKaratzas 直接批准。

- 暂无高价值评论线程

风险与影响

- 风险：回归风险：移除 [nvidia/Kimi-K2.6-NVFP4](#) 意味着该模型的 NVFP4 MoE 测试暂时跳过。若后续有修改影响该模型，CI 无法捕获，需在重新添加后回归。性能影响：模块级 fixture 可能延长首次运行时间，但避免了每函数重复下载，整体测试时间应减少。兼容性：无，仅修改测试文件。
- 影响：CI 稳定性：直接修复 AMD CI 的 Kernels Quantization 测试组 6 小时超时问题，使其能按预期快速完成。测试覆盖：暂时降低了对 [nvidia/Kimi-K2.6-NVFP4](#) 模型 NVFP4 量化路径的测试覆盖，但保持了对 [nvidia/Qwen3-30B-A3B-NVFP4](#) 的覆盖。团队：AMD 硬件 CI 受益，开发者可更快获得测试反馈。

- 风险标记：CI 超时修复，测试覆盖暂降

关联脉络

- PR #44667 （假设：引入 Kimi-K2.6-NVFP4 的 PR）：该 PR 被指引入导致超时的模型下载，但具体 PR 未在上下文中提及，此处为猜测