

PR #46773 完整报告

vllm-project/vllm

[ModelRunner V2] Fix whisper test

合并时间: 2026-06-26 13:10

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/46773>

执行摘要

- 一句话: 修复 Whisper 测试因 EncoderCache 缺少 `__len__` 方法失败
- 推荐动作: 推荐管理员快速合并。这是一个小但必要的修复, 解决了测试阻塞问题, 且经过审查和少量改动。值得关注的细节是: 条件判断的 `is None` 改法体现了 vLLM 对 Python 代码质量的一贯追求。

功能与动机

Whisper 测试用例 (`tests/models/multimodal/generation/test_whisper.py#L363`) 中使用了 `len(encoder_cache)` 来检查缓存是否为空, 而 `EncoderCache` 类未实现 `__len__` 方法, 导致测试运行时抛出异常。PR body 明确引用该测试行作为上下文。

实现拆解

1. 在 `EncoderCache` 类中新增 `__len__` 方法 (文件: `vllm/v1/worker/gpu/mm/encoder_cache.py`, +3 行)。新增的 `__len__` 方法返回 `self.encoder_outputs` 字典的长度, 使得 `len(encoder_cache)` 能够正常工作, 直接满足测试需求。
2. 修复 `mm_pruning.py` 中的条件判断 (文件: `vllm/v1/worker/gpu/model_states/mm_pruning.py`, +3/-3 行)。将 `not rope_state`、`not encoder_cache`、`not model_config.multimodal_config` 分别改为 `rope_state is None`、`encoder_cache is None`、`model_config.multimodal_config is None`。这是更安全的 `None` 检查方式, 避免因这些变量具有 `__bool__` 方法或为其他假值而产生误判。虽然这部分改动与 Whisper 测试无直接关系, 但有助于提升代码健壮性。

关键文件:

- `vllm/v1/worker/gpu/mm/encoder_cache.py` (模块 编码器缓存; 类别 source; 类型 core-logic; 符号 `len`): 新增 `__len__` 方法, 直接修复测试失败的根本原因。
- `vllm/v1/worker/gpu/model_states/mm_pruning.py` (模块 多模态剪枝; 类别 source; 类型 data-contract): 修改了条件判断中的 `None` 检查方式, 提升代码健壮性。

关键符号: `EncoderCache.len`, `maybe_create_mm_pruner`

关键源码片段

`vllm/v1/worker/gpu/mm/encoder_cache.py`

新增 `__len__` 方法，直接修复测试失败的根本原因。

```
# SPDX-License-Identifier: Apache-2.0
# SPDX-FileCopyrightText: Copyright contributors to the vLLM project
import torch

from vllm.multimodal.inputs import MultiModalFeatureSpec

class EncoderCache:
    def __init__(self):
        # req_id -> MM features
        self.mm_features: dict[str, list[MultiModalFeatureSpec]] = {}
        # MM hash -> encoder outputs
        self.encoder_outputs: dict[str, torch.Tensor] = {}

    # 新增 __len__ 方法，返回 encoder_outputs 字典的长度
    # 这使得 len(encoder_cache) 能够正常工作，修复 Whisper 测试中的调用
    def __len__(self) -> int:
        return len(self.encoder_outputs)

    def add_request(
        self, req_id: str, mm_features: list[MultiModalFeatureSpec]
    ) -> None:
        self.mm_features[req_id] = mm_features

    def remove_request(self, req_id: str) -> None:
        self.mm_features.pop(req_id, None)

    def reset_mm_cache(self) -> None:
        """
        Clear the multi-modal cache that was used during profiling,
        but no longer needed during inference.
        """
        # TODO: Implement MM budget for encoder dummy run
        pass

    def reset_encoder_cache(self) -> None:
        """Clear the GPU-side encoder cache storing vision embeddings.

        This should be called when model weights are updated to ensure
        stale embeddings computed with old weights are not reused.
        """
        self.encoder_outputs.clear()

    def free_encoder_cache(self, mm_hash: str) -> None:
        self.encoder_outputs.pop(mm_hash, None)
```

评论区精华

仅有一条人类审核评论：yewentao256 审核并批准（LGTM），无其他讨论。claude[bot] 自动评论因 PR 来自 fork 而跳过自动化审查。

- 暂无高价值评论线程

风险与影响

- 风险：风险极低。改动量小（+6/-3 行），只涉及新增 1 个方法和修改条件判断语法。
__len__ 方法直接委托给 dict 的 len，不会引入新错误。条件判断的修改虽然改变了语义（从隐式的布尔检查改为显式的 None 检查），但更符合 Python 最佳实践，且 rope_state、encoder_cache、model_config.multimodal_config 在现有代码中预计只有 None / 非 None 两种状态，因此回归风险极低。
- 影响：直接影响：Whisper 测试用例现在可以正常运行。间接影响：任何其他依赖 len(encoder_cache) 的代码都能正常工作。mm_pruning.py 的改动提升了代码可读性和正确性，但不会改变现有行为。影响范围仅限于 ModelRunner V2 的多模态相关模块。
- 风险标记：暂无

关联脉络

- 暂无明显关联 PR