

PR #46213 完整报告

vllm-project/vllm

[Bugfix][Multimodal] Fix Qwen3-Omni use_audio_in_video with mixed image/video inputs

合并时间: 2026-07-17 16:31

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/46213>

执行摘要

- 一句话: 修复 Qwen3-Omni 在 use_audio_in_video 时混合图像 / 视频输入崩溃与模态分类错误
- 推荐动作: 该 PR 值得精读, 尤其适合关注多模态输入处理和 API 设计的开发者。核心设计决策 (使用 Tensor 属性传递元数据而非修改函数签名) 提供了一个简洁的模式, 可在其他需要为数据附加临时信息的场景中复用。check_interleaved_audio_video 的重写展示了如何正确处理含边界 token 的多跨度交错检测。

功能与动机

用户在使用 Qwen3-Omni 的 use_audio_in_video=True 功能时, 遇到单视频 + 图片输入 (V+I, I+V) 导致 EngineCore 崩溃, 同时发现混合图像 / 视频 / 音频输入时模态判断错误、多视频输入因边界 token 被误分类为非交错。PR body 详细描述了三个具体 bug 并提供了测试矩阵, 旨在使所有组合通过。

实现拆解

1. 重新实现交错检测逻辑: 在 qwen2_5_omni_thinker.py 中重写 check_interleaved_audio_video(), 从全局密集范围检查改为遍历每个连续的 V/A 跨度, 每个跨度独立判断是否包含视频和音频 token 且范围重叠, 从而避免多视频边界 token 导致的误判。
2. 引入显式模态元数据机制: 在 vllm/multimodal/utils.py 中新增三个工具函数 —— set_mm_embedding_modality()、copy_mm_embedding_modality()、get_mm_embedding_modalities(), 通过向 embedding tensor 动态添加 .modality 属性来记录其真实模态 (video/audio/image), 后续 merge 时直接读取该属性而非通过 token 计数猜测。
3. 在编码器与模型主路径中传播模态信息: 修改 vllm/v1/worker/gpu_model_runner.py 的 _gather_mm_embeddings(), 在收集每个 mm_embeds_item 后立即调用 set_mm_embedding_modality() 设置模态; m-rope 位置重计算后, 通过 copy_mm_embedding_modality() 将模态传递到新张量。同步修改 mm_pruning.py 的 strip() 和 recompute() 以保持模态属性。
4. 调整 Qwen3-Omni 模型入口: 在 qwen3_omni_moe_thinker.py 的 embed_input_ids() 中添加图像 token 掩码 (is_image), 并修正 is_vision = is_video | is_image, 使交错分支正确识别所有视觉占位符。同时更新 merge_interleaved_embeddings() 调用, 移除不再需

要的 `num_video` 和 `num_audio` 参数。

5. 配套测试：更新 `test_qwen2_5_omni_embed.py`，新增 `test_multi_video_with_boundary_tokens`（验证多视频场景）、`test_image_and_video_mixed`（验证图像 + 视频混合）、`test_missing_modality_raises`（验证缺少模态属性时抛出），并调整现有用例以使用 `_mm_embed` 辅助函数（内部调用 `set_mm_embedding_modality`）。新增 `test_gather_preserves_mixed_modalities` 测试编码器收集时保留混合模态。

关键文件：

- `vllm/multimodal/utils.py`（模块 多模态；类别 source；类型 core-logic；符号 `set_mm_embedding_modality`, `copy_mm_embedding_modality`, `get_mm_embedding_modalities`）：新增三个核心工具函数 `set_mm_embedding_modality`, `copy_mm_embedding_modality`, `get_mm_embedding_modalities`，通过在 `embedding tensor` 上附加 `.modality` 属性来传递模态信息，替代了原本需贯穿整个调用链的参数列表传递方式，使变更影响面大幅缩小。
- `vllm/model_executor/models/qwen2_5_omni_thinker.py`（模块 模型层；类别 source；类型 data-contract；符号 `_merge_embedding_group`）：核心 bugfix 文件：重写 `check_interleaved_audio_video()` 为逐个连续跨度检测，解决多视频边界误判；重写 `merge_interleaved_embeddings()` 使用显式模态属性合并 `embedding`，并引入内部函数 `_merge_embedding_group` 提升可维护性。
- `tests/models/multimodal/processing/test_qwen2_5_omni_embed.py`（模块 测试；类别 test；类型 test-coverage；符号 `_mm_embed`, `test_multi_video_with_boundary_tokens`, `test_image_and_video_mixed`, `test_missing_modality_raises`）：新增关键测试用例覆盖多视频边界、图像 + 视频混合、模态缺失等场景，并引入 `_mm_embed` 辅助函数以支持修饰后的 `embedding` 构造，与源码改动紧密对应。

关键符号： `check_interleaved_audio_video`, `merge_interleaved_embeddings`, `_merge_embedding_group`, `set_mm_embedding_modality`, `copy_mm_embedding_modality`, `get_mm_embedding_modalities`, `_gather_mm_embeddings`, `embed_input_ids` (`qwen3_omni_moe_thinker`), `strip` (`mm_pruning`), `recompute` (`mm_pruning`)

关键源码片段

`vllm/multimodal/utils.py`

新增三个核心工具函数 `set_mm_embedding_modality`、`copy_mm_embedding_modality`、`get_mm_embedding_modalities`，通过在 `embedding tensor` 上附加 `.modality` 属性来传递模态信息，替代了原本需贯穿整个调用链的参数列表传递方式，使变更影响面大幅缩小。

```
def set_mm_embedding_modality(embed: "torch.Tensor", modality: str) -> "torch.Tensor":  
    """Attach modality metadata to a gathered multimodal embedding tensor.
```

```
  
    Used by interleaved Omni merge paths that need to group embeddings by  
    modality without threading a parallel modalities list through  
    ``embed_input_ids``.
```

```

"""
embed.modality = modality # type: ignore[attr-defined]
return embed

def copy_mm_embedding_modality(
    src: "torch.Tensor", dst: "torch.Tensor"
) -> "torch.Tensor":
    """Copy ``modality`` from ``src`` onto ``dst`` if present."""
    modality = getattr(src, "modality", None)
    if modality is not None:
        dst.modality = modality # type: ignore[attr-defined]
    return dst

def get_mm_embedding_modalities(
    multimodal_embeddings: Sequence["torch.Tensor"],
) -> list[str]:
    """Collect per-embedding modalities previously set on the tensors."""
    modalities: list[str] = []
    for i, emb in enumerate(multimodal_embeddings):
        modality = getattr(emb, "modality", None)
        if modality is None:
            raise ValueError(
                f"Missing modality on multimodal embedding at index {i}. "
                "Encoder gather must set embed.modality before interleaved "
                "audio-in-video merge."
            )
        modalities.append(modality)
    return modalities

```

评论区精华

关键设计讨论：使用 Tensor 属性而非额外返回值

在 review 中，@Isotr0py 建议通过 `mm_embeds_item.modality = mm_feature.modality` 的方式将模态信息直接附加到 embedding tensor 上，而非通过返回额外列表再逐层传递，因为后者需要刷新所有模型的接口。@gty111 表示同意并实际实现了这一方案（回复 "Done"），最终代码采用了属性方式，使变更范围显著缩小。

兼容性关注

@gty111 在 `gpu_model_runner.py` 的 `_gather_mm_embeddings` 处提问：“请验证其他 `_gather_mm_embeddings` 调用是否与此更改兼容。”最终调整后，所有调用点均保持兼容。

CI 结果确认

@gty111 验证了修复效果并指出失败的 CI 与 PR 无关，由 @Isotr0py 确认后合并。

- 使用 Tensor 属性替代额外返回参数传递模态信息 (design): 采用属性方式，简化变更范围。

- 验证其他调用点的兼容性 (correctness): 经过实际修改和测试, 确认所有 `_gather_mm_embeddings` 调用点兼容。
- CI 失败与 PR 无关 (other): CI 失败不阻碍合并。

风险与影响

- 风险:
 - 核心路径变更: 修改了 `_gather_mm_embeddings`、`check_interleaved_audio_video` 等关键多模态处理函数, 可能影响其他使用相同基础设施的模型 (如 Qwen2.5-VL、Qwen3-VL)。但采用的属性传递方式是非侵入式的, 旧路径不受影响。
 - 属性依赖: 动态添加的 `.modality` 属性依赖 Python 对象的 `__dict__`, 在 tensor 视图或切片上可能丢失。当前使用场景明确 (仅在 merge 前使用), 且提供了 `copy_mm_embedding_modality` 用于新张量复制。
 - 缺少回归测试: 虽然新增了单元测试, 但没有覆盖完整端到端流程的回归测试。PR body 中提供了手动测试矩阵, 但未纳入 CI 自动运行。
 - 性能影响: 属性设置开销可忽略, 无显著性能风险。
- 影响:
 - 用户: 修复了 `use_audio_in_video=True` 时图像 + 视频混合输入崩溃, 直接改善 Qwen3-Omni 用户体验; 多视频场景和混合模态分类错误不再出现, 结果正确。
 - 系统: 修改了多模态数据流中的模态记录方式, 为未来其他模型 (如 Gemma、Mistral) 支持类似功能提供了可复用工具函数。
 - 团队: 代码可读性和可维护性提升, 属性方案避免了复杂的参数传递链。影响范围限于 Qwen3-Omni 及 Qwen2.5-Omni 模型路径。
 - 风险标记: 核心路径变更, 属性依赖 Tensor monkey-patching, 缺少端到端回归测试, 仅影响 Qwen3-Omni 路径

关联脉络

- 暂无明显关联 PR