

PR #49484 完整报告

vllm-project/vllm

Fix GLM-4.1V video placeholder token ID handling.

合并时间: 2026-07-25 03:35

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/49484>

执行摘要

- 一句话: 修复 GLM-4.1V 视频占位符 token ID 错误
- 推荐动作: 建议精读核心改动, 尤其是 `_call_hf_processor` 中的 token 映射逻辑和 `_get_video_frame_embed_token_id` 的抽取, 这是管理 token ID 一致性的良好实践。

功能与动机

ROCm 多模态 CI 中测试 `test_custom_inputs_models[glm4_1v-video-test_case7]` 失败。问题根源是视频帧占位符使用了 `video_token_id` 而非 `image_token_id`, 导致生成的描述与 HuggingFace 不一致 (例如 "book on floor" 对比 "laptop on bed")。

实现拆解

在 `vllm/model_executor/models/glm4_1v.py` 中新增 `_get_video_frame_embed_token_id` 方法, 根据处理器类型返回正确的 token ID: 对于 `Glm4vProcessor` 或 `TRANSFORMERS_WITH_GA` 时返回 `image_token_id`, 否则返回 `video_token_id`。

用 `_get_video_frame_embed_token_id` 替换原有的硬编码条件, 确保 GLM-4.1V 视频帧占位符使用 `image_token_id`。

仅在 `_get_video_frame_embed_token_id` 返回 `image_token_id` 时 (即 GLM-4.1V), 对视频块内的 `image_token_id` 执行替换为 `video_token_id` 的映射, 以兼容 HuggingFace 的文本回馈路径。同时新增反向映射, 将处理后的 `video_token_id` 恢复为 `image_token_id`, 以保持一致的最终 prompt。

将 `get_image_replacement_glm4v` 中的 token ID 从 `video_token_id` 改为 `image_token_id`, 简化条件逻辑。

`tests/models/multimodal/processing/test_glm4_1v.py` 中 `test_processor_override` 改为直接使用 `hf_processor.image_token_id` 进行计数, 移除多余的 tokenizer 变量。

关键文件:

- `vllm/model_executor/models/glm4_1v.py` (模块 模型执行; 类别 source; 类型 data-contract; 符号 `_get_video_frame_embed_token_id`, `get_image_replacement_glm4v`, `get_image_replacement`, `get_video_replacement_glm4v`): 核心源码文件, 包含所有逻辑变更: 新增 `_get_video_frame_embed_token_id` 方法, 修改 `_construct_video_placeholder` 和

`_call_hf_processor` 以及 `get_image_replacement_glm4v` 等函数。

- `tests/models/multimodal/processing/test_glm4_1v.py` (模块 GLM 1v; 类别 test; 类型 test-coverage) : 测试文件, 确保视频占位符使用正确的 token ID 计数。

关键符号: `_get_video_frame_embed_token_id`, `get_image_replacement_glm4v`, `_construct_video_placeholder`, `_call_hf_processor`

关键源码片段

`vllm/model_executor/models/glm4_1v.py`

核心源码文件, 包含所有逻辑变更: 新增 `_get_video_frame_embed_token_id` 方法, 修改 `_construct_video_placeholder` 和 `_call_hf_processor` 以及 `get_image_replacement_glm4v` 等函数。

```
def _get_video_frame_embed_token_id(self, hf_processor: object) -> int:    # GLM-4.1V
    (Glm4vProcessor) 和带 GA 的 GLM 模型使用 image_token_id 作为帧嵌入    if
    isinstance(hf_processor, Glm4vProcessor) or TRANSFORMERS_WITH_GA:        return
    hf_processor.image_token_id    # 其他 GLM 模型使用 video_token_id    return
    hf_processor.video_token_id 这个新增方法将之前分散的条件逻辑集中管理, 使得代码意图清
    晰。在 _call_hf_processor 中, 条件映射逻辑如下:    # 仅在帧嵌入使用 image_token_id 时,
    才需要在 HF 处理后交换 token swap_video_frame_tokens = frame_embed_token_id ==
    processor.image_token_id if swap_video_frame_tokens:    input_ids[input_ids ==
    processor.image_token_id] = processor.video_token_id # ... 后续处理 ... # 最后再将
    video_token_id 映射回 image_token_id if swap_video_frame_tokens:
    input_ids[input_ids == processor.video_token_id] = processor.image_token_id
```

评论区精华

此次 review 没有实质讨论。AndreasKaratzas 指出该 PR 与 #49220 内容相同, 因之前合并导致 main 损坏而迁移至此, 且已被两位 committer 批准。Claude bot 因 PR 来自 fork 而跳过自动审查。

- 暂无高价值评论线程

风险与影响

- 风险: 低风险。变更集中在 GLM-4.1V 的视频占位符生成路径, 通过提取方法和条件映射隔离了影响。但仍需注意以下风险:
 - 回归风险: `_call_hf_processor` 中的条件 token 映射可能影响混合图文输入。
 - 兼容风险: 依赖 `TRANSFORMERS_WITH_GA` 标志, 若 transformers 版本变化可能改变行为。
 - 影响: 影响面较小, 仅涉及 GLM-4.1V 模型的多模态处理。修复后, 视频帧占位符与 HuggingFace 一致, 生成描述准确性得到提升。对系统性能无影响。
- 风险标记: 依赖 transformers 版本标志 `TRANSFORMERS_WITH_GA`

关联脉络

- PR #49220 Fix GLM-4.1V video placeholder token ID handling.: 该 PR 与当前 PR 内容完全相同，因之前合并导致 main 损坏，迁移至此 PR。