

PR #49477 完整报告

vllm-project/vllm

[Perf] Defer MM embeds loading off the event loop

合并时间: 2026-07-23 21:58

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/49477>

执行摘要

- 一句话: 将 MM embeds 加载异步化以释放事件循环
- 推荐动作: 该 PR 设计简洁、测试充分、性能提升明显, 值得合并。建议其他开发者注意: 未来新增 embedding 类型时应同步提供异步方法以避免类似问题。

功能与动机

PR body 明确指出: `With --enable-mm-embeds, loading multi-MB embedding tensors blocks the API server's event loop, stalling all concurrent requests.` 作者通过微基准测试展示了 128MB embedding 加载时事件循环阻塞从 110ms 降至 54ms。

实现拆解

1. 新增异步方法: 在 `vllm/multimodal/media/connector.py` 中新增 `fetch_image_embedding_async` 和 `fetch_audio_embedding_async`, 通过 `loop.run_in_executor(global_thread_pool, ...)` 将同步的 `load_base64` 调用委托给线程池, 避免阻塞事件循环。
2. 重构请求构建器: 在 `vllm/entrypoints/chat_utils.py` 中, `parse_image_embeds` 和 `parse_audio_embeds` 不再直接调用同步方法并立即传给 `tracker`, 而是注册一个异步包装器 `_image_embeds_with_uuid_async` / `_audio_embeds_with_uuid_async`, 由 `tracker` 在合适时机 `await`。
3. 并发优化: 对于 dict 输入 (多 key), 使用 `asyncio.gather` 并发获取所有 embedding, 然后重新组合为 dict。
4. 保留同步路径: 原有的 `fetch_image_embedding` 和 `fetch_audio_embedding` 方法保持不变, 仅新增异步版本, 无向后兼容风险。

关键文件:

- `vllm/entrypoints/chat_utils.py` (模块 请求前端; 类别 source; 类型 core-logic; 符号 `_image_embeds_with_uuid_async`, `_audio_embeds_with_uuid_async`): 核心重构所在, 将同步加载替换为异步包装器, 并利用 `asyncio.gather` 并发处理 dict 输入。
- `vllm/multimodal/media/connector.py` (模块 多模态; 类别 source; 类型 core-logic; 符号 `fetch_image_embedding_async`, `fetch_audio_embedding_async`): 新增的两个异步方法 `fetch_image_embedding_async` 和 `fetch_audio_embedding_async`, 通过 `run_in_executor` 将同步 IO 卸载到线程池。

关键符号: `_image_embeds_with_uuid_async`, `_audio_embeds_with_uuid_async`,
`fetch_image_embedding_async`, `fetch_audio_embedding_async`

关键源码片段

`vllm/entrypoints/chat_utils.py`

核心重构所在, 将同步加载替换为异步包装器, 并利用 `asyncio.gather` 并发处理 dict 输入。

```
# vllm/entrypoints/chat_utils.py —— 新增的异步包装器
async def _image_embeds_with_uuid_async(
    self,
    image_embeds: str | dict[str, str] | None,
    uuid: str | None,
):
    # 将原本在 parse_image_embeds 中同步执行的加载操作移到此处,
    # 由 tracker 在事件循环空闲时 await, 避免阻塞事件循环
    if isinstance(image_embeds, dict):
        # 使用 asyncio.gather 并发加载所有 key 对应的 embedding
        tensors = await asyncio.gather(
            *(
                self._connector.fetch_image_embedding_async(v)
                for v in image_embeds.values()
            )
        )
        embeds = dict(zip(image_embeds, tensors))
    elif isinstance(image_embeds, str):
        embeds = await self._connector.fetch_image_embedding_async(image_embeds)
    else:
        embeds = None
    return embeds, uuid # 返回 (item, uuid) 元组满足 tracker 契约
```

`vllm/multimodal/media/connector.py`

新增的两个异步方法 `fetch_image_embedding_async` 和 `fetch_audio_embedding_async`, 通过 `run_in_executor` 将同步 IO 卸载到线程池。

```
# vllm/multimodal/media/connector.py —— 新增的异步方法
async def fetch_image_embedding_async(
    self,
    data: str,
) -> torch.Tensor:
    """
    Asynchronously load image embedding from a URL.
    """
    image_embedding_io = ImageEmbeddingMediaIO()
    loop = asyncio.get_running_loop()
    # 将同步的 load_base64 调用提交到全局线程池,
    # 避免阻塞事件循环; 返回一个 awaitable
    return await loop.run_in_executor(
        global_thread_pool, image_embedding_io.load_base64, "", data
```

)

```
async def fetch_audio_embedding_async(self, data: str) -> torch.Tensor:
    audio_embedding_io = AudioEmbeddingMediaIO()
    loop = asyncio.get_running_loop()
    return await loop.run_in_executor(
        global_thread_pool, audio_embedding_io.load_base64, "", data
    )
```

评论区精华

唯一实质性 review 来自 DarkLight1337，在 [vllm/entrypoints/chat_utils.py](#) 上建议对 dict 输入使用 `asyncio.gather` 以并发获取所有 embedding。作者 Guan-Ming 随即采纳该建议并更新代码（第二次 commit）。此外无其他技术讨论。

- 使用 `asyncio.gather` 并发获取 dict embedding (performance): 作者采纳建议，将 dict 分支改为 `asyncio.gather` 并发加载，提升性能。

风险与影响

- 风险：变更仅涉及新增异步方法和重构调用点，原有同步方法未改动，回归风险低。主要风险在于：若 `global_thread_pool` 被过度占用，可能影响其他同步操作；但该线程池已在其他模态（如图像、音频）中广泛使用，风险可控。测试覆盖方面，未新增专门测试，但现有 `tests/entrypoints/unit_tests/test_chat_utils.py` 通过（63 passed）。
- 影响：用户层面：使用 `--enable-mm-embeds` 时，多模态请求的并发性能显著提升，事件循环阻塞时间从加载耗时降低到仅 GIL 持有时间。系统层面：减少了 API 服务器的请求排队延迟，提升整体吞吐。团队层面：代码符合现有异步模态加载模式，易于维护。影响范围限定于启用了 mm embeds 的场景。
- 风险标记：核心路径变更，缺少测试覆盖（无新增测试）

关联脉络

- PR #49317（未在已知列表中，但 PR body 提及）：PR body 说明本 PR 并非 duplicate，issue #49317 覆盖不同阶段