

PR #26082 完整报告

sgl-project/sglang

perf: eliminate CUDA syncs in VLM embed path

合并时间: 2026-06-12 13:15

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/26082>

执行摘要

- 一句话: 消除 VLM 嵌入路径中的 CUDA 同步
- 推荐动作: 值得精读: 作者识别并消除了典型 CUDA 同步模式, 使用 `masked_scatter_`、CPU 偏移计算和预计算哈希三种互补技术。设计讨论展示了务实的工程权衡 (推迟不常见优化)。

功能与动机

Profiling identified three unnecessary CUDA stream synchronizations in `mm_utils.py` that stall the GPU pipeline: `torch.where(mask)`, `mask.sum().item()`, and SHA256 re-hash on `ShmPointerMMData`. These syncs serialize GPU execution and degrade throughput for multi-image and video workloads.

实现拆解

1. 用 `masked_scatter_` 替换 `torch.where`: 在 `embed_mm_inputs` 中新增 `scatter` 辅助函数, 使用 `dest.masked_scatter` 替代 `torch.where` + 索引赋值, 避免隐式同步。
2. CPU 端多模态令牌计数: 新增 `_count_mm_tokens_from_offsets` 函数, 从 `items_offset_list` 等 CPU 侧元数据直接计算 True 数量, 避免 `mask.sum().item()` 造成的数据传输。当 EVS 修改 `input_ids` 时仍回退到原方法。
3. 预计算哈希转发: `ShmPointerMMData.init__` 接收 `precomputed_hash`; `hash_feature` 中优先返回缓存哈希, 避免耗时的 SHA256 计算。同时通过 `__getstate__`/`__setstate__` 保证序列化往返完整性。
4. 空列表保护: 在 `hash_feature` 的 `isinstance` 检查前添加 `len(f) > 0` 判断, 防止空特征列表触发 `IndexError`。
5. 测试配置修复: 为 `TestKimiVLServer` 添加 `--mem-fraction-static=0.40`, 避免大图像请求导致的 OOM。

关键文件:

- `python/sglang/srt/managers/mm_utils.py` (模块 多模态嵌入; 类别 `source`; 类型 `core-logic`; 符号 `_scatter`, `ShmPointerMMData.init`, `ShmPointerMMData.getstate`, `ShmPointerMMData.setstate`): 核心变更文件, 包含所有同步消除逻辑 (`_scatter`、预计算哈希转发、空列表保护) 和序列化扩展。

- test/registered/vlm/test_vision_openai_server_a.py（模块 测试配置；类别 test；类型 test-coverage）：修复 Kimi-VL 测试的 OOM 问题，确保 CI 通过。

关键符号：_scatter, hash_feature, ShmPointerMMData.init, ShmPointerMMData.getstate, ShmPointerMMData.setstate

关键源码片段

python/sglang/srt/managers/mm_utils.py

核心变更文件，包含所有同步消除逻辑（_scatter、预计算哈希转发、空列表保护）和序列化扩展。

```
# embed_mm_inputs 内新增 _scatter 避免 torch.where 的 CUDA 同步
# masked_scatter_ 在行主序掩码下的行为与 torch.where 排序索引相同
# 无需 materialize 索引张量，从而消除同步

def _scatter(dest, mask, src):
    dest.masked_scatter_(mask.expand_as(dest), src.to(dest.device, dest.dtype))

for i, modality, embedding, mask in zip(
    range(len(embeddings)), modalities, embeddings, masks
):
    if embedding is None or mask is None:
        continue
    _scatter(input_embeds, mask, embedding)
    if use_deepstack.get(modality, None):
        _scatter(input_deepstack_embeds, mask, deepstack_embeddings[i])

# hash_feature 优先使用预计算哈希，避免 SHA256 重新计算
# 同时确保空列表不会触发 IndexError

def hash_feature(f):
    if isinstance(f, list):
        if len(f) > 0 and isinstance(f[0], ShmPointerMMData):
            # 直接对底层张量哈希，暂未利用单个元素的预计算哈希
            return tensor_hash([x.tensor for x in f])
        if len(f) > 0 and isinstance(f[0], torch.Tensor):
            return tensor_hash(f)
        return data_hash(tuple(flatten_nested_list(f)))
    ...
    elif isinstance(f, ShmPointerMMData):
        if f.precomputed_hash is not None:
            return f.precomputed_hash # 直接返回缓存值
        return tensor_hash([f.tensor])

# ShmPointerMMData 序列化时携带 precomputed_hash，保证跨进程传递
class ShmPointerMMData:
    def __init__(self, tensor: torch.Tensor, precomputed_hash: Optional[int] = None):
        ...
```

```
self.precomputed_hash = precomputed_hash

def __getstate__(self):
    return {"shm_name": self.shm_name, "shape": self.shape,
            "dtype": self.dtype, "precomputed_hash": self.precomputed_hash}

def __setstate__(self, state):
    ...
    self.precomputed_hash = state.get("precomputed_hash")
```

评论区精华

- Gemini 自动审查建议对列表中的 ShmPointerMMData（如视频帧）也利用预计算哈希，避免全部 rehash。
- ShangmingCai起初提议用 `data_hash(tuple(precomputed_hashes))`，但后来认为该情况过于边缘，决定不在本 PR 包含，避免过度工程。
- ShangmingCai指出 `_wrap_tensor_or_list` 可能需要传递哈希到列表元素，但代价太大，可能为死代码。
- ShangmingCai报告 lint 报错 `get_chunked_prefill_embedding_legacy` 不存在，可能为暂存问题。
- Optimizing hashing for list of ShmPointerMMData objects (performance): Deferred: reviewer considered it over-protected and rare edge case, not included in this PR to avoid complexity.
- Completeness of precomputed hash forwarding through `_wrap_tensor_or_list` (design): Acknowledged but deliberately excluded as potential dead code; may be addressed separately.
- Lint error about missing `_legacy` function (correctness): Not fully resolved; assumed transient as CI passed after fixing OOM and retrying unrelated failures.

风险与影响

- 风险：
 - `masked_scatter_` 语义等价性：仅当掩码为行主序连续时才等价于 `torch.where`；当前掩码形状为 `[seq_len, 1]` 且展开为 `[seq_len, hidden_dim]`，填充顺序与排序索引一致，风险低。
 - 预计算哈希依赖：若 `precomputed_hash` 在序列化中丢失（未正确设置或旧数据），`hash_feature` 通过 `__setstate__` 的 `.get()` 降级为重新计算，但需保证 sender 侧确实计算了哈希。
 - 测试覆盖有限：仅调整了单行配置，未新增针对同步消除的专项测试。
- 影响：
 - 用户影响：多模态推理吞吐量提升，尤其是多图像和视频场景。
 - 系统影响：减少 GPU 同步，提升流水线利用率。

- 团队影响：核心 `mm_utils.py` 路径更干净，但需留意 `masked_scatter_` 在极端掩码模式下的行为差异。
- 风险标记：核心路径同步消除，预计算哈希序列化依赖，测试覆盖仅 CI 修复

关联脉络

- 暂无明显关联 PR