

PR #36386 完整报告

sgl-project/sglang

[HiCache] Heal the storage existence cache on a hybrid prefetch discard

合并时间: 2026-08-28 16:33

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/36386>

执行摘要

- 一句话: 修复混合预取丢弃后存在缓存未愈合的问题
- 推荐动作: 该 PR 值得精读, 因为它解决了一个特定的缓存一致性 bug, 展示了如何通过局部失效来保持缓存正确性。设计上通过计算 `keep_pages` 精确控制失效范围, 避免过度失效, 体现了对性能的考量。建议关注其计算逻辑的边界条件, 并考虑补充单元测试以覆盖该路径。

功能与动机

在混合布局 (SWA/Mamba 组件、混合 stacks) 中, 存储预取是全部或全无的: 如果任何辅助池未能取回完整前缀, `_check_hybrid_prefetch_result` 会丢弃整个结果。但为该哈希链记录的 KV 存在性信念未被丢弃, 导致存在性缓存持续声称该 span 存在而辅助页缺失, 后续预取信任这些陈旧的正向结果并以相同方式失败。循环中没有逻辑能闭合这个漏洞。

实现拆解

1. 在 `python/sglang/srt/mem_cache/unified_radix_cache.py` 的 `_check_hybrid_prefetch_result` 方法中, 在混合预取丢弃分支内新增 `keep_pages` 计算逻辑。
2. 计算方式: 初始值设为 `completed_tokens // self.page_size` (KV 完成的页数), 然后对每个辅助传输进行修正: 如果 `transfer.keys` 是 `None` (传输无键), 则 `keep_pages` 置为 0; 否则如果传输的键数少于 `transfer.keys` 的长度 (即传输不足), 则 `keep_pages` 取 `min(keep_pages, max(0, len(hash_value) - len(transfer.keys)))`, 因为辅助传输以链的尾部页为键。
3. 调用 `self.storage_existence_cache.invalidate_beyond(PoolName.KV, hash_value, keep_pages=keep_pages)` 从第一个未服务的页面开始丢弃 KV 信念, 使下一个插入通过一次 FULL 检查重写该 span, 从而恢复缺失的辅助页。
4. 在现有的丢弃警告日志中报告 `keep_pages` 以便观察愈合情况。
5. 由于缓存模式不填充存在性缓存, `invalidate_beyond` 在缓存模式下为 no-op, 因此缓存模式行为不变。
6. 测试与基准: 本 PR 未添加测试或基准, 因为不涉及模型前向路径变更, 仅涉及缓存信念失效。

关键文件:

- python/sclang/srt/mem_cache/unified_radix_cache.py (模块 缓存层; 类别 source; 类型 core-logic; 符号 _check_hybrid_prefetch_result) : 核心方法
_check_hybrid_prefetch_result 的丢弃分支中加入存在性缓存失效逻辑, 修复残留陈旧正向信念导致重复预取失败的问题。

关键符号: _check_hybrid_prefetch_result

关键源码片段

python/sclang/srt/mem_cache/unified_radix_cache.py

核心方法 `_check_hybrid_prefetch_result` 的丢弃分支中加入存在性缓存失效逻辑, 修复残留陈旧正向信念导致重复预取失败的问题。

```
# 来自 python/sclang/srt/mem_cache/unified_radix_cache.py 的 _check_hybrid_prefetch_result 方法
```

```
if pool_transfers and not all_succeeded:
```

```
    # 关键修复: 丢弃预取结果时, 同步愈合存储存在性缓存。
```

```
    # 计算所有池实际提供了的前导页数 keep_pages,
```

```
    # 从第一个未服务的页开始丢弃 KV 信念,
```

```
    # 这样下一次插入会通过一次 FULL 检查重写该 span, 恢复缺失的辅助页。
```

```
    keep_pages = completed_tokens // self.page_size
```

```
    for transfer, count in zip(pool_transfers, pool_hit_pages):
```

```
        if transfer.keys is None:
```

```
            # 传输没有任何 key, 说明辅助池完全未覆盖, 保留 0 页。
```

```
            keep_pages = 0
```

```
        elif count < len(transfer.keys):
```

```
            # 辅助传输以链的尾部页为 key,
```

```
            # 如果实际命中页数少于请求的 key 数, 则说明尾部有缺失,
```

```
            # 保留页数应减去缺失的尾部页数。
```

```
            keep_pages = min(
```

```
                keep_pages, max(0, len(hash_value) - len(transfer.keys))
```

```
            )
```

```
    # 只无效化从 keep_pages 之后的 KV 信念, 保留已一致的前导部分。
```

```
    self.storage_existence_cache.invalidate_beyond(
```

```
        PoolName.KV, hash_value, keep_pages=keep_pages
```

```
    )
```

```
    ...
```

```
    logger.warning(
```

```
        "HiCache hybrid prefetch discarded req=%s completed=%d requested=%d "
```

```
        "kv_beliefs_kept_pages=%d", # 在告警中报告保留页数, 便于观测。
```

```
        req_id,
```

```
        completed_tokens,
```

```
        expected_tokens,
```

```
        keep_pages,
```

```
    )
```

评论区精华

该 PR 没有 review 评论, 只有一个 /tag-and-rerun-ci 命令。因此没有实质性的讨论线程。

- 暂无高价值评论线程

风险与影响

- 风险:

1. 该改动位于混合预取丢弃路径，属于核心缓存管理逻辑的一部分，但仅增加了一个 $O(\text{pages})$ 的循环和一次缓存失效调用，对热路径无影响。
2. `keep_pages` 的计算依赖 `transfer.keys` 和 `pool_hit_pages` 的语义，若不一致可能导致过早失效（过度失效）或失效不足（残留陈旧信念）。但代码注释和逻辑已考虑辅助传输以尾部页为键的特性，风险较低。
3. 日志格式变更（增加 `kv_beliefs_kept_pages` 字段）不会影响功能，但依赖日志解析的工具可能需要更新。
4. 未新增单元测试，回归风险由现有测试覆盖，但该特定逻辑未被测试覆盖。 - 影响：该 PR 影响 HiCache 混合布局（SWA/Mamba）下的预取行为，修复了重复无效预取导致的性能问题。对使用纯 KV 布局的用户无影响（缓存模式为 `no-op`）。影响范围较小，仅触及一个文件，改动量小，但能显著减少混合布局下因陈旧存在性信念导致的重复预取失败。
- 风险标记：核心路径变更，缺少测试覆盖

关联脉络

- PR #35931 [HiCache] Reject load-back specs that claim nodes pinned by an in-flight load-back: 同样涉及 HiCache 的调度器崩溃问题，与缓存节点 pin 相关，可能与本次修复的重复预取问题有一定重叠。
- PR #36227 [HiCache] Retry L3 storage prefetch after a missed attempt: 该 PR 处理了预取失败后的重试逻辑，与本 PR 修复的重复预取失败问题在功能上相关。
- PR #36705 [HiCache] Stop populating host-pool mmmaps twice (-13% allocation time): 也涉及 HiCache 的预取性能优化，与本 PR 在性能改进方向上一致。