

PR #29792 完整报告

sgl-project/sglang

Fix Mamba track-boundary bookkeeping under overlap scheduling

合并时间: 2026-08-12 00:36

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/29792>

执行摘要

- 一句话: 修复重叠调度下 Mamba track 边界记账竞态
- 推荐动作: 值得精读。PR body 对竞态的推演 (尤其是“wait_stream 为什么只是掩盖问题”以及 mamba_next_track_idx 双重语义的拆解) 是理解 sglang overlap scheduling 与 HiCache 前缀缓存交互的绝佳材料; producer/consumer 指针分离与“per-batch 快照 vs 共享可变状态”的取舍可作为并发状态管理的范例。建议重点阅读 batch_result_processor.py 与 memory_pool.py 的配套改动, 以及 hanming-lu 的 follow-up 统一提交。

功能与动机

PR body 指出: Mamba 每次 decode forward 用 batch 本地 `seq_lens_cpu` 快照构建 `mamba_track_mask`, 但结果处理侧却用共享可变的 `Req.kv_committed_len` 判定是否跨过 track 边界; 在 overlap scheduling 下, 处理 batch A 结果时该请求可能已被 batch B 的 `prepare_for_decode()` 推进, 导致 Mamba 后处理跑错批次, lazy extra-buffer 模式下会错误地推进 / 释放 ping-pong 槽并记录不匹配的 `mamba_last_track_seqlen`, 后续 cache-hit replay 载入与缓存前缀不一致的 Mamba 状态, 造成 KL 发散。同时 `mamba_next_track_idx` 被复用于“下一次 forward 要写的槽”和“最近一次合法状态所在槽”两种含义, 重叠调度下二者不一致, cache 插入可能保留预分配的未来槽。作者强调 wait_stream 只是同步层面的掩盖, 正确修复是让结果处理使用与 forward 相同的 per-batch 快照, 并显式保留已提交的 Mamba 状态槽。

实现拆解

1. 定位根因 (结果处理侧与 forward 不同源): batch_result_processor.py 的 `_mamba_check_track_boundary` 原先直接用 `req.kv_committed_len % interval == 0` 判定边界, 改为 `committed_len = req.kv_committed_len - lookahead`, 其中 `lookahead = req.decode_batch_idx - batch.mamba_decode_batch_idx_cpu[i]`, 使结果处理与 forward 的 `seq_lens_cpu` 快照语义一致; lookahead in (0, 1) 断言继续约束重叠窗口。
2. 引入 producer/consumer 双指针模型: 在 Req 上新增 `mamba_last_track_idx` (`schedule_batch.py`), 表示“最近一次已提交合法 Mamba 状态所在槽”, 与 producer 指针 `mamba_next_track_idx` 解耦; memory_pool.py 的 `get_mamba_ping_pong_keep_idx` 统一返回 `mamba_last_track_idx` (不再按 lazy 模式分支), `_alloc_ping_pong_buffer` 完成初始化 (lazy 为 0、非 lazy 为 other(0)), `free_mamba_cache`、`reset_for_retract`、

streaming_session.SessionSlot 的 save/restore、MLX auxiliary_state.py 与 model_runner.py 全部同步保存 / 恢复 / 清空该字段，避免悬空指针。

3. 重构结果处理时序逻辑：_handle_finish_state_updated_req 区分 completed_mamba_boundary（本 batch forward 实际完成跨界）与 known_mamba_boundary（提交跨界），lazy 模式用 completed_mamba_boundary 决定是否执行 _mamba_prefix_cache_update，非 lazy 用 known_mamba_boundary；非 lazy 完成跨界时用 batch 快照 mamba_track_buffer_indices[i] 与 kv_committed_len - lookahead 回填 mamba_last_track_idx / mamba_last_track_seqlen；新增 lazy + finished + lookahead == 1 且 next == last 时置 mamba_lazy_is_insert = False 的守卫，防止 in-flight 下一 batch 覆盖同一 keep 槽后仍错误插入缓存。
mamba_lazy_post_decode_at_boundary(req, batch, track_idx) 改为接收 batch 快照槽位，将两个指针收敛到 track_idx 并释放另一槽。
4. 数据契约配套：ScheduleBatch 新增 mamba_track_buffer_indices（每 batch 逻辑槽位 0/1 快照），在 set_mamba_track_indices_from_reqs 中与物理 mamba_track_indices 一同生成，并随 filter_batch / merge_batch / copy 传播；prefill 与 prepare_for_decode 路径同步维护新指针。
5. 测试与验证配套：test_unified_radix_cache_kl_mamba.py 将 HiCache IO 后端由 direct 切到 kernel、内存布局由 page_first_direct 切到 page_first，覆盖新 IO 后端下的 KL 一致性；PR body 给出 Qwen3-Next-80B-A3B-Instruct-FP8 + TP4 + overlap 下的 GSM8K (0.953) 与 MMLU (0.854) 精度结果。提交历史共 21 个 commit，经历多次 main 合并与 radix_cache/unified_radix_tree、test_qwen3_next_models.py 的反复重跑；合并前 hanming-lu 以 follow-up commit 将 keep-idx 语义在 lazy/non-lazy 全路径统一到 mamba_last_track_idx。

关键文件：

- python/sglang/srt/managers/scheduler_components/batch_result_processor.py（模块 结果处理；类别 source；类型 core-logic；符号 _handle_finish_state_updated_req, _mamba_prefix_cache_update, _mamba_check_track_boundary, mamba_lazy_post_decode_at_boundary）：核心修复点：结果处理侧 Mamba 边界判定从共享可变的 req.kv_committed_len 改为还原 lookahead 后的 per-batch 视图，并重构 _handle_finish_state_updated_req、_mamba_prefix_cache_update、mamba_lazy_post_decode_at_boundary 的更新时序。
- python/sglang/srt/mem_cache/memory_pool.py（模块 内存池；类别 source；类型 core-logic；符号 get_mamba_ping_pong_keep_idx, _alloc_ping_pong_buffer, free_mamba_cache）：get_mamba_ping_pong_keep_idx 从按 lazy 模式推导改为直接读 mamba_last_track_idx, _alloc_ping_pong_buffer 初始化新指针，free_mamba_cache 清理新指针，是修复 cache 插入保留错误槽位的核心配套。
- python/sglang/srt/managers/schedule_batch.py（模块 调度批次；类别 source；类型 core-logic；符号 Req, ScheduleBatch, set_mamba_track_indices_from_reqs）：新增 Req.mamba_last_track_idx 字段与 ScheduleBatch.mamba_track_buffer_indices 每 batch 逻辑槽位快照，set_mamba_track_indices_from_reqs 生成快照并随 filter/merge/copy 传播，是整个修复的数据契约基础。

- python/sglang/srt/session/streaming_session.py (模块 会话缓存; 类别 source; 类型 core-logic; 符号 SessionSlot, save_from_req, restore_to_req) : SessionSlot 需在流式会话跨轮保存 / 恢复 mamba_last_track_idx, 否则会话恢复的请求会回退到竞态的 mamba_next_track_idx 判定 keep 槽。
- python/sglang/srt/hardware_backend/mlx/kv_cache/auxiliary_state.py (模块 MLX 缓存; 类别 source; 类型 core-logic; 符号 free_mamba_cache, cleanup_after_caching_req) : MLX 后端的 free_mamba_cache 与 cleanup_after_caching_req 需同步清理 mamba_last_track_idx, 避免 MLX auxiliary-state 缓存路径出现悬空 keep 指针。
- python/sglang/srt/hardware_backend/mlx/model_runner.py (模块 MLX 执行; 类别 source; 类型 data-contract; 符号 _store_tracked_auxiliary_state) : MLX _store_tracked_auxiliary_state 新建 track buffer 时需初始化 mamba_last_track_idx = 0, 保证新请求生命周期内 keep-idx 有定义。
- test/registered/radix_cache/unified_radix_tree/test_unified_radix_cache_kl_mamba.py (模块 回归测试; 类别 test; 类型 test-coverage) : Mamba KL 回归测试的 HiCache 后端配置从 direct/page_first_direct 切到 kernel/page_first, 是本 PR 修复场景的验收测试, 也是 CI 反复重跑的核心对象。

关键符号: _handle_finish_state_updated_req, _mamba_prefix_cache_update, _mamba_check_track_boundary, mamba_lazy_post_decode_at_boundary, _mamba_lazy_spec_update, get_mamba_ping_pong_keep_idx, _alloc_ping_pong_buffer, set_mamba_track_indices_from_reqs, donate_mamba_ping_pong_slot, save_from_req, restore_to_req

关键源码片段

python/sglang/srt/mem_cache/memory_pool.py

get_mamba_ping_pong_keep_idx 从按 lazy 模式推导改为直接读 mamba_last_track_idx, _alloc_ping_pong_buffer 初始化新指针, free_mamba_cache 清理新指针, 是修复 cache 插入保留错误槽位的核心配套。

memory_pool.py —— keep 槽位语义统一收敛到 mamba_last_track_idx

```
def get_mamba_ping_pong_keep_idx(self, req: Req) -> int:
    """返回持有最近一次已提交 tracked 状态的 ping-pong 槽位。"""
    # 无论 lazy 与否统一使用 consumer 指针:
    # 它只在结果处理确认写入完成后才更新, 不受重叠调度中
    # producer 指针 (mamba_next_track_idx) 提前推进的影响
    return req.mamba_last_track_idx
```

```
def _alloc_ping_pong_buffer(self, req: Req):
    """为新请求分配 ping-pong track buffer。
```

```
    lazy 模式分配 1 个槽、第二个置 -1 (边界处按需分配);
    普通模式一次性分配全部槽。初始时 last_track_idx 指向首个
    有效状态所在槽: lazy 为 0, 非 lazy 为 other(0)。
    """
```

```

n = (
    1
    if self.enable_mamba_extra_buffer_lazy
    else self.mamba_ping_pong_track_buffer_size
)
slots = self.mamba_allocator.alloc(n)
assert slots is not None, (
    "Not enough space for mamba ping pong idx, "
    "try to increase --mamba-full-memory-ratio."
)
buf = torch.full(
    (self.mamba_ping_pong_track_buffer_size,),
    -1,
    dtype=slots.dtype,
    device=slots.device,
)
buf[:n] = slots
req.mamba_ping_pong_track_buffer = buf
req.mamba_next_track_idx = 0
# 初始化 consumer 指针，保证请求生命周期内 keep-idx 始终有定义
req.mamba_last_track_idx = (
    0
    if self.enable_mamba_extra_buffer_lazy
    else self.get_mamba_ping_pong_other_idx(0)
)

```

python/sglang/srt/managers/schedule_batch.py

新增 `Req.mamba_last_track_idx` 字段与 `ScheduleBatch.mamba_track_buffer_indices` 每 batch 逻辑槽位快照，`set_mamba_track_indices_from_reqs` 生成快照并随 filter/merge/copy 传播，是整个修复的数据契约基础。

schedule_batch.py —— 每 batch 快照 Mamba 逻辑槽位（数据契约）

```
def set_mamba_track_indices_from_reqs(batch, track_positions=None):
```

```
    """从 req 构建 mamba_track_indices（权威来源）。
```

```
    track_positions 为可选的每请求乒乓位置覆盖（lazy spec 的
    track plan）。除 GPU forward 使用的物理槽 ID（mamba_track_indices）
    外，同时把逻辑位置 0/1 快照到 mamba_track_buffer_indices：
    结果处理侧在重叠调度下据此还原本次 forward 真正写入的槽位。
    """
```

```
    req_to_token_pool = batch.req_to_token_pool
    all_buffers = req_to_token_pool.req_index_to_mamba_ping_pong_track_buffer_mapping[
        batch.req_pool_indices
    ] # (bs, ping_pong_size), int64, on device
    if track_positions is None:
        track_positions = [
            req.mamba_next_track_idx if req.mamba_next_track_idx is not None else 0
            for req in batch.reqs

```

```

]
# 逻辑位置快照：结果处理时 req 级指针可能已被下一 batch 推进
batch.mamba_track_buffer_indices = list(track_positions)
idx = (
    torch.tensor(track_positions, dtype=torch.int64, pin_memory=True)
    .unsqueeze(1)
    .to(device=all_buffers.device, non_blocking=True)
)
batch.mamba_track_indices = (
    torch.gather(all_buffers, 1, idx).squeeze(1).to(torch.int64)
)

```

评论区精华

hanming-lu (issue 评论) : `mamba_last_track_idx` gives me the feeling that it's only for getting around the `extra_buffer_lazy` kl test failure. I feel it's more like a workaround instead of a proper fix?

作者在 PR body 中系统性论证这是 ownership/indexing 语义 bug 的根治而不是 workaround : `wait_stream` 只是改变重叠时序、让错误路径观察到已完成状态从而掩盖问题，正确做法是显式保留已提交状态槽。最终 hanming-lu 以 follow-up commit 将 `keep_idx` 语义在 `lazy/non-lazy` 全路径统一到 `mamba_last_track_idx`，等于接受该设计并进一步收敛。

hanming-lu (`schedule_batch.py` 2089 行) : what's the difference between `mamba_track_buffer_indices` and `mamba_track_indices`?

作者回答: `mamba_track_buffer_indices` 是 per-batch 快照的逻辑 ping-pong 位置 (0/1)，`mamba_track_indices` 是对应的物理 Mamba pool 槽 ID (供 GPU forward 使用)；结果处理时必须用前者，因为 req 级指针在重叠调度下可能已被下一 batch 推进。ispobock 随后要求补充解释性注释，已添加。

YazhiGao (`batch_result_processor.py`) : would appreciate we do a per batch torch tensor -> int get, .item can be pretty costly on weak cpu plat like grace

性能提醒被采纳：最终 `_mamba_check_track_boundary` 改为纯整数运算 `req.kv_committed_len - lookahead`，避免逐请求 tensor 到 int 的同步。

yizhang2077: why do we need remove it (指 `mamba_lazy_post_decode_at_boundary` 中的 `req.finished()` 分支)

作者解释：该函数现在只处理已完成的 batch A；若 lazy prealloc 失败，A 会在原地写完 checkpoint，该 checkpoint 有效仍应插入缓存，而 in-flight batch B 写同一槽的危险场景已由 `lookahead == 1 && known_mamba_boundary && next == last` 守卫单独处理，保留旧分支会让 alloc-fail 场景下所有请求错过缓存。

hanming-lu (中间版本 `_mamba_has_inflight_write_to_keep_slot`) : why are all three check necessary?

该辅助函数在后续演进中被 `_handle_finish_state_updated_req` 中更简洁的三条件守卫取代，代码复杂度下降。

- `mamba_last_track_idx` 是否是绕过 KL 测试的 workaround (design): 作者用 batch A/B 的逐步推演说服评审, hanming-lu 后续以 follow-up commit 将 `keep_idx` 语义在 `lazy/non-lazy` 全路径统一到 `mamba_last_track_idx`, 质疑转化为更彻底的一致性收敛。
- `mamba_track_buffer_indices` 与 `mamba_track_indices` 的差异 (question): 作者说明前者是 per-batch 逻辑槽位快照、后者是 GPU forward 用的物理 slot ID, 并补充了代码注释。
- 逐个 `.item()` 在弱 CPU 平台上的开销 (performance): 最终实现改为纯整数运算 `req.kv_committed_len - lookahead`, 避免逐请求 tensor 到 int 同步。
- 删除 `mamba_lazy_post_decode_at_boundary` 中的 `req.finished()` 分支 (correctness): 删除分支, in-flight 写同一槽的危险由 `lookahead == 1 && known_mamba_boundary && next == last` 守卫单独处理。
- `get_mamba_ping_pong_keep_idx` 的 if-else 表达 (style): 后续演进中简化为直接返回 `req.mamba_last_track_idx`, 在统一消费指针语义后不再需要后备分支。
- `mamba_track_buffer_indices` 缺少解释性注释 (documentation): 作者按要求补充代码注释, 解释快照的用途与重叠调度背景。

风险与影响

- 风险: `_handle_finish_state_updated_req` 是每个 decode 结果必经的热路径, 改动影响所有启用 Mamba + HiCache + overlap scheduling 的部署; `lookahead` in (0, 1) 断言硬性限制重叠窗口, 若未来调度器扩大重叠深度会直接触发 `assert` (属防御性保护而非优雅降级)。新增字段涉及多处所有权转移路径 (session 保存 / 恢复、retract 重置、MLX 释放、内存池 free), 本次已全覆盖, 但任何后续新增后端若漏同步 `mamba_last_track_idx`, 会导致 keep 槽悬空或 use-after-free。测试后端从 `direct/page_first_direct` 切到 `kernel/page_first`, 旧 IO 后端路径在该回归测试中不再被覆盖。此外 CI 多次出现与本次改动无直接关系的失败 (`test_qwen3_next_models.py`、`kl_dsv4_pp`、`kl_mimo`、`kl_nightly`), 存在 flaky 噪音, 合入后需留意 nightly 稳定性。
- 影响: 用户侧: 修复 `--enable-mamba-extra-buffer-lazy + overlap scheduling + HiCache` 下 Mamba 状态缓存错位导致的 KL 发散与精度不稳定, 作者给出 GSM8K 0.953 / MMLU 0.854 的精度佐证。系统侧: 不引入全局 forward-stream wait, 避免同步性能损失; 仅增加每 batch 一个 `mamba_track_buffer_indices` 快照与少量指针维护, 热路径开销可忽略。团队侧: 确立了 Mamba ping-pong 状态的 producer/consumer 双指针语义契约, 后续所有涉及 `mamba_next_track_idx` / `mamba_last_track_idx` 的改动 (session、MLX、radix cache 插入) 都必须遵守; `get_mamba_ping_pong_keep_idx` 的语义从“模式相关推导”收敛为“直接读 consumer 指针”, 降低了认知负担。
- 风险标记: 核心调度路径变更, 重叠调度时序竞态, 新增字段需全路径同步, 测试后端切换, CI 存在 flaky 噪音

关联脉络

- PR #32208 `O(1) slot allocation in ReqToTokenPool.alloc()`: 同样改动 `python/sglang/srt/mem_cache/memory_pool.py` 并涉及 Req 级槽位所有权与调度热路径语义, 与本 PR 的 Mamba ping-pong 槽位所有权修复同一演进线。

- PR #34356 Add bit-exact hicache logprob-consistency test: HiCache + Mamba 的 KL 一致性测试线，与本 PR 修复的 KL 发散现象直接相关，共同保障 HiCache 缓存状态位精确性。
- PR #34405 Fix flaky decode cache-hit check in Inkling test: HiCache 测试 flaky 修复，与本 PR 反复重跑 radix_cache/unified_radix_tree 与模型 e2e 测试的稳定性诉求同源。