

# PR #28434 完整报告

sgl-project/sglang

[HiCache]Support hybrid pool staged H2D kernel

合并时间: 2026-06-19 09:48

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/28434>

## 执行摘要

- 一句话: 支持混合池 staged H2D 写回 JIT 内核
- 推荐动作: 建议阅读本 PR 以了解 HiCache 如何从单池 staged 扩展到混合池。关键设计点包括:
  1. 将 staged module 与主 module 分离, 降低编译依赖;
  2. 通过 `can_use_write_back_jit` 标志解耦写回路径选择;
  3. CPU 模拟测试框架 (mock JIT 函数) 值得复用。

对于 Mamba 精度回归, 建议后续 PR 跟进确定性能差异根因或在 Mamba 池中回退到非 staged 路径。

## 功能与动机

之前 page-first 布局已默认启用, 但混合池 (如 DeepSeek V4 MLA、Mamba 等) 的写回仍无法使用 staged JIT 内核, 且在某些场景下会因 `move_indices` 触发 `RuntimeError: Destination indices must be a CUDA tensor`。本 PR 旨在统一混合池的写回路径, 消除该错误。

## 实现拆解

1. 独立 staged 模块加载: 在 `python/sglang/jit_kernel/hicache.py` 中将原有的 `_jit_hicache_module` 分离, 新增 `_jit_hicache_staged_module`, 仅编译 `staged_write_back.cuh`, 不再依赖 `hicache.cuh` 和 `layout.cuh`。新增 `can_use_write_back_jit_kernel` 函数, 对 `element_size` 的最低对齐要求从 128 字节降为 16 字节, 使更多配置可用。
2. 统一初始化暂存缓冲区: 在 `python/sglang/srt/mem_cache/memory_pool_host.py` 中, 为 `MHATokenToKVPoolHost`、`MLATokenToKVPoolHost`、`MambaPoolHost`、`DeepSeekV4PagedHostPool`、`DeepSeekV4StateHostPool`、`DSASearcherPoolHost` 等池的 `_init_write_back_staging_buffers` 方法增加 `can_use_write_back_jit` 检测。当 layout 为 `page_first` 且 JIT kernel 可用时, 分配 staging 缓冲区并置标志为 `True`; 否则保持 `False`。该标志后续用于写回时的分支选择。
3. 写回路径切换: 在各池的 `backup_from_device_all_layer` 方法中, 条件从 `self.can_use_jit` 改为 `self.can_use_write_back_jit`。当新标志为 `True` 时, 调用 staged 版本的 JIT 函数 (如 `jit_transfer_hicache_all_layer_staged_if_pf`、

`jit_transfer_hicache_all_layer_mla_staged_lf_pf`)，这些函数内部改为使用 `_jit_hicache_staged_module`。

4. 控制器适配：在 `python/sglang/srt/managers/cache_controller.py` 和 `python/sglang/srt/mem_cache/hybrid_cache/hybrid_cache_controller.py` 的 `start_writing` 中，判断是否使用 staged 写回时增加 `getattr(self.mem_pool_host, "can_use_write_back_jit", False)` 条件。该条件使 host indices 可以保持 CPU 端，避免不必要的 GPU 拷贝和 `move_indices` 调用。
5. 测试配套：新增 `test/registered/unit/mem_cache/test_hicache_staged_write_back_dispatch.py` (926 行)，通过 mock 所有 JIT 函数为 CPU 端模拟实现，覆盖 MHA、MLA、Mamba、DeepSeek V4、DSA Indexer 等池的 staged 写回和加载路径。同时对 `test/registered/jit/test_hicache.py` 进行调整，改用 `can_use_write_back_jit_kernel` 检测，并在测试中验证写回后再加载的正确性。其他 KL 回归测试配置统一修改为 `kernel + page_first`。

关键文件：

- `python/sglang/srt/mem_cache/memory_pool_host.py` (模块 缓存层；类别 source；类型 core-logic；符号 `_init_write_back_staging_buffers`, `init`)：核心修改文件，为所有混合池类型添加 `can_use_write_back_jit` 标志和 staged 暂存缓冲区，并切换写回路径的条件。
- `test/registered/unit/mem_cache/test_hicache_staged_write_back_dispatch.py` (模块 单元测试；类别 test；类型 test-coverage；符号 `_indices`, `_ptr_key_from_layers`, `_ptr_key_from_tensor`, `_cpu_staged_lf_pf_copy`)：新增的 926 行单元测试，通过 CPU 模拟 mock 所有 JIT 函数，覆盖全部池类型的 staged 写回和加载分发逻辑。
- `python/sglang/jit_kernel/hicache.py` (模块 JIT 内核；类别 source；类型 core-logic；符号 `_jit_hicache_staged_module`, `can_use_write_back_jit_kernel`)：将 staged 模块从主 hicache 模块中分离，新增 `_jit_hicache_staged_module` 和 `can_use_write_back_jit_kernel`，降低了编译耦合，并为更多配置提供支持。
- `test/registered/jit/test_hicache.py` (模块 JIT 测试；类别 test；类型 test-coverage；符号 `test_hicache_page_first_staged_write_back_mha_staged_only_alignment`, `test_hicache_page_first_staged_write_back_mla_staged_only_alignment`)：调整现有 JIT 测试，使用 `can_use_write_back_jit_kernel` 检测替代 `host_pool.can_use_jit`，并增加写回后再加载的验证步骤，确保 staged 路径双向正确。
- `python/sglang/srt/managers/cache_controller.py` (模块 调度器；类别 source；类型 entrypoint)：在 `start_writing` 中增加 `can_use_write_back_jit` 检测，使 staged 写回时保持 host indices 在 CPU，避免 `move_indices` 错误。
- `python/sglang/srt/mem_cache/hybrid_cache/hybrid_cache_controller.py` (模块 混合缓存；类别 source；类型 entrypoint)：与 `cache_controller.py` 相同的改动，适配混合缓存控制器。
- `python/sglang/srt/mem_cache/hybrid_cache/hybrid_pool_assembler.py` (模块 混合池组装；类别 source；类型 core-logic)：调整配置键，与 `can_use_write_back_jit` 对齐。
- `test/registered/radix_cache/unified_radix_tree/test_unified_radix_cache_kl_mamba.py` (模块 KL 测试；类别 test；类型 test-coverage)：调整测试配置为 `kernel+page_first`，以匹配新的写回路径；该测试因本 PR 出现精度回归 (maxKL 上升)。

关键符号: `_init_write_back_staging_buffers`, `backup_from_device_all_layer`, `_jit_hicache_staged_module`, `can_use_write_back_jit_kernel`, `start_writing`, `_cpu_staged_lf_pf_copy`, `_cpu_staged_mha_lf_pf_copy`, `_cpu_jit_one_layer_mha_copy`, `_cpu_jit_one_layer_mla_copy`, `_cpu_per_layer_pf_lf_copy`

## 关键源码片段

`python/sglang/srt/mem_cache/memory_pool_host.py`

核心修改文件，为所有混合池类型添加 `can_use_write_back_jit` 标志和 `staged` 暂存缓冲区，并切换写回路径的条件。

```
# python/sglang/srt/mem_cache/memory_pool_host.py
```

```
# 以 MHATokenToKVPoolHost 为例，展示 _init_write_back_staging_buffers 和写回路径变更
```

```
class MHATokenToKVPoolHost(BaseHostKVCache):
```

```
    def _init_write_back_staging_buffers(self):
```

```
        # 初始化暂存缓冲区属性，默认关闭
```

```
        self.staging_page_capacity = 0
```

```
        self.staging_token_capacity = 0
```

```
        self.staging_k_buffer = None
```

```
        self.staging_v_buffer = None
```

```
        self.can_use_write_back_jit = False
```

```
        # 只有 page_first 布局且非 NPU/XPU/MPS 才尝试启用 staged 写回
```

```
        if self.layout != "page_first" or (_is_npu or _is_xpu or _is_mps):
```

```
            return
```

```
        # 探测是否能加载 staged JIT kernel，要求 element_size % 16 == 0
```

```
        self.can_use_write_back_jit = _is_cuda and can_use_write_back_jit_kernel(
            element_size=self.element_dim * self.dtype.itemsize,
        )
```

```
        if not self.can_use_write_back_jit:
```

```
            return
```

```
        # 分配 staged buffer，大小为一个 chunk (_WRITE_BACK_STAGING_PAGE_CHUNK) 乘以
        page_size
```

```
        self.staging_page_capacity = min(
            self.page_num, _WRITE_BACK_STAGING_PAGE_CHUNK
        )
```

```
        self.staging_token_capacity = self.staging_page_capacity * self.page_size
```

```
        self.staging_k_buffer = torch.empty(
            (self.staging_token_capacity, self.layer_num, self.element_dim),
            dtype=self.dtype,
            device=self.device,
        )
```

```
        self.staging_v_buffer = torch.empty_like(self.staging_k_buffer)
```

```
    def backup_from_device_all_layer(
```

```

        self, device_pool, host_indices, device_indices, backend="kernel"
    ):
        # ... 其他分支 ...
        elif self.layout == "page_first":
            # 根据 can_use_write_back_jit 选择 JIT staged 还是通用路径
            if self.can_use_write_back_jit:
                # 调用 staged 写回 JIT 函数, 使用暂存缓冲区
                jit_transfer_hicache_all_layer_staged_lf_pf(
                    k_ptr_src=device_pool.k_data_ptrs,
                    v_ptr_src=device_pool.v_data_ptrs,
                    src_indices=device_indices,
                    dst_indices=host_indices,
                    staging_k=self.staging_k_buffer,
                    staging_v=self.staging_v_buffer,
                    dst_k=self.k_buffer,
                    dst_v=self.v_buffer,
                    page_size=self.page_size,
                )
            else:
                # 回退到逐层 copy 的通用路径
                self._backup_from_device_per_layer_general(
                    device_pool, host_indices, device_indices
                )

```

## test/registered/unit/mem\_cache/test\_hicache\_staged\_write\_back\_dispatch.py

新增的 926 行单元测试, 通过 CPU 模拟 mock 所有 JIT 函数, 覆盖全部池类型的 staged 写回和加载分发逻辑。

```

# test/registered/unit/mem_cache/test_hicache_staged_write_back_dispatch.py
# CPU 模拟 staged 写回拷贝函数示例

```

```

def _cpu_staged_lf_pf_copy(
    src_registry,
    *,
    ptr_src,
    src_indices,
    dst_indices,
    dst,
    **_,
):
    # "src_registry" 是一个字典, key 为 ptr 元组, value 为层列表
    src_layers = src_registry[_ptr_key_from_tensor(ptr_src)]
    src_indices = src_indices.to(dtype=torch.int64, device="cpu")
    dst_indices = dst_indices.to(dtype=torch.int64, device="cpu")
    for layer_id, src in enumerate(src_layers):
        dst[dst_indices, layer_id] = src[src_indices]

```

```

class TestHiCacheStagedWriteBackDispatch(unittest.TestCase):

```

```

def _patched_transfers(self, src_registry=None):
    # 返回一组 mock.patch, 将真实的 JIT 函数替换为 CPU 模拟
    staged_side_effect = None
    if src_registry is not None:
        staged_side_effect = lambda **kwargs: _cpu_staged_lf_pf_copy(
            src_registry, **kwargs
        )
    return (
        mock.patch(
            f"{MEMORY_POOL_HOST_MODULE}.jit_transfer_hicache_all_layer_staged_lf_pf",
            side_effect=staged_side_effect,
        ),
        # ... 其他 mock ...
    )

```

### python/sglang/jit\_kernel/hicache.py

将 staged 模块从主 hicache 模块中分离，新增 `_jit_hicache_staged_module` 和 `can_use_write_back_jit_kernel`，降低了编译耦合，并为更多配置提供支持。

```

# python/sglang/jit_kernel/hicache.py
# 新增的 staged 模块加载函数和探测函数

```

```

@cache_once
def _jit_hicache_staged_module(
    *, element_size: int, unroll: int, block_quota: int
) -> Module:
    args = make_cpp_args(
        element_size,
        unroll,
        block_quota,
        1024, # num_threads, 为模板兼容保留
    )
    return load_jit(
        "hicache_staged", # 独立模块名，避免与主模块冲突
        *args,
        cuda_files=[
            "kvcacheio/staged_write_back.cuh", # 只编译 staged 相关 kernel
        ],
        cuda_wrappers=[
            ("launch_all_lf_pf_staged",
             f"&HiCacheStagedWriteBackKernel<{args}>::run_all_lf_pf_staged"),
            ("launch_all_mla_lf_pf_staged",
             f"&HiCacheStagedWriteBackKernel<{args}>::run_all_mla_lf_pf_staged"),
        ],
    )

```

```

def can_use_write_back_jit_kernel(
    *,

```

```

    element_size: int,
    unroll: int | None = None,
    block_quota: int | None = None,
) -> bool:
    logger = logging.getLogger(__name__)
    # 对齐要求从 128 字节降为 16 字节，使更多 element_size 可用
    if element_size % 16 != 0:
        logger.warning(f"unsupported {element_size=} for staged JIT kernel")
        return False
    try:
        unroll = unroll or _default_unroll(element_size)
        block_quota = block_quota or DEFAULT_BLOCK_QUOTA
        _jit_hicache_staged_module(
            element_size=element_size,
            unroll=unroll,
            block_quota=block_quota,
        )
        return True
    except Exception as e:
        logger.warning(f"Failed to load staged JIT HiCache kernel: {e}")
        return False

```

## 评论区精华

- gemini-code-assist[bot] 建议优化：在 `MambaPoolHost._init_write_back_staging_buffers` 中，建议直接使用预计算的状态元素大小和 dtype 计算 `element_size`，而非通过 helper 方法索引已分配的缓冲区；在 `DSAIndexerPoolHost._init_write_back_staging_buffers` 中，指出 `self.indexer_page_stride_size` 已包含 `indexer_dtype.itemsize` 乘法，再次乘以会导致冗余。这些建议尚未明确是否被采纳，`memory_pool_host.py` 最终提交中上述点可能已优化或保持原样。
- ispobock 报告的 mamba KL regression: `test_unified_radix_cache_kl_mamba.py` 在 PR 合并后 maxKL 从 0.003-0.0045 上升到 0.007-0.015，经二分确认为此 PR 引入。该回归集中在 Mamba 的 conv/temporal state 写回路径。小组多次 rerun 后 CI 最终通过，但未明确修复或回归容忍度调整。这暗示 staged 写回对 Mamba 存在精度差异，当前通过调整测试阈值或使用确定性 kernel 解决，但未来可能需进一步校准。
- 两位 reviewer 批准：xiezhq-hermann 和 yuan-luo 均无条件批准，表明主要设计决策已被团队认同。
  - gemini-code-assist 优化建议：MambaPoolHost 和 DSAIndexerPoolHost 中 `element_size` 的计算 (design): 建议未被明确确认是否采纳，但最终提交中相关点可能已优化或保持原样（提交记录未显示额外修改）。
  - ispobock 报告 mamba KL 回归经二分确认为本 PR 引入 (correctness): 经过多次 rerun 后 CI 最终通过，但未明确修复或调整测试阈值。可能通过非确定性 kernel 或随机 seed 容忍。团队接受了该回归水平。
  - 代码结构：独立 staged 模块 vs 合并在主模块中 (design): 决定采用独立模块，通过 `@cache_once` 保证只加载一次。

## 风险与影响

- 风险：
  - Mamba 精度回归：KL 测试显示 Mamba 状态写回后 maxKL 有显著上升 (0.003 → 0.015)，尽管 CI 最终通过，但精度损失可能是非确定性的，可能影响长推理场景的缓存命中质量。建议后续跟踪长序列的端到端准确率。
  - 核心路径变更：写回路径是 HiCache 的关键数据流，`can_use_jit` 改为 `can_use_write_back_jit` 后，若 JIT 编译失败或 staging 未正确初始化，会回退到非 JIT 路径，但不会崩溃。`getattr(..., 'can_use_write_back_jit', False)` 的 fallback 也保证了向后兼容。
  - JIT 编译时间：独立 staged 模块后，两个模块分别加载，可能增加启动时的编译总时间，但均为一次缓存。
  - 测试覆盖：新增的单元测试覆盖了所有池类型和边界 (`page_count=1/63/64/65/128/129`)，使用 CPU 模拟 JIT，未直接调用真实 CUDA kernel，可能漏掉 GPU 端的数据竞争或同步问题。建议后续添加 GPU 端集成测试。
  - 性能影响：staged 写回相比直接写回多了一次 staging 缓冲区的复制，但通过分块 (`staging_page_capacity`) 限制显存消耗，大页面数下吞吐应高于非 JIT 路径。无端到端 benchmark，性能风险可控。
- 影响：
  - 用户影响：使用 HiCache 且 `layout=page_first` 的用户自动获得混合池的 staged 写回优化，消除潜在的 `RuntimeError`，无配置变更。Mamba 用户可能感知细微精度差异（参见风险）。
  - 系统影响：写回路径统一，后续新增池类型只需实现 `_init_write_back_staging_buffers` 即可复用 staged 机制。独立 staged 模块降低了编译耦合，便于单独优化。
  - 团队影响：增加了约 1300 行源码，其中测试占 1000+ 行，保持了较好的测试密度。引入 `can_use_write_back_jit` 概念，与 `can_use_jit` 并存的命名可能造成短期混淆，但注释和代码设计清晰。
  - 影响程度：中高。涉及 HiCache 核心写回路径，四个控制器 (`cache_controller.py`、`hybrid_cache_controller.py`) 均修改，回归风险集中在上游 Mamba 精度。
  - 风险标记：核心路径变更，KL 精度回归风险，混合池首次支持，测试覆盖充分，JIT 编译时间增加

## 关联脉络

- PR #21631 [HiCache] page-first default config for staged write-back: 本 PR 是在 #21631 引入的 page-first staged write-back 基础上的扩展，将支持范围从标准 MHA/MLA 池扩展到所有混合池。
- PR #28755 Cap SWA pool sizing with chunk cache: 近期对 SWA 池大小上限的调整与 HiCache 的内存管理相关，但与本 PR 无直接功能依赖。