

PR #39726 完整报告

vllm-project/vllm

[SimpleCPUOffloadConnector]: Add support for reset_cache()

合并时间: 2026-06-18 10:47

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/39726>

执行摘要

- 一句话: 为 SimpleCPUOffloadConnector 添加 reset_cache()
- 推荐动作: 建议 PR 审核者重点验证与 dedup 修复 (PR#41289) 的配合, 确认 `_in_flight_store_gpu_blocks` 在 reset 中正确清理。对于 KV offload 领域的开发者, 本 PR 的 'abandoned transfer' 设计模式值得学习, 用于在异步系统中安全重置状态。一般读者可跳过实现细节, 只关注决策过程。

功能与动机

PR 描述指出, `Scheduler.reset_prefix_cache()` 在 in-flight DMA 传输时失败, 因为 `ref_cnt>0` 的块仍被使用。本 PR 旨在添加 `reset_cache()` 支持, 通过设计让所有改动局限在 SimpleCPUOffloadConnector 内, 避免与 worker 同步, 如 pr body 所述: 'These changes purposely are done this way to allow zero changes outside of the SimpleCPUOffloadConnector and not need to sync with the worker.'

实现拆解

1. 引入 abandoned 状态字典: 在 `SimpleCPUOffloadScheduler.__init__()` 中添加 `_abandoned_store_event_to_blocks` 和 `_abandoned_reqs_to_load` 字典, 用于在 `reset()` 时将未完成的 store 和 load 事件移入, 保证引用不会被立即释放, 直到对应传输完成。
2. 实现 `reset()` 方法: 在 `manager.py` 中新增 `reset()` 方法。该方法将当前所有 `_store_event_to_blocks` 和 `_reqs_to_load` 移入 abandoned 字典, 清除调度器队列、CPU 前缀缓存, 重置光标 (lazy 模式)。如果 abandoned 字典非空则返回 `False`, 所有 pending 传输完成后返回 `True`。要求 `_gpu_block_pool` 必须已绑定, 否则断言失败。
3. 调整传输完成处理: 修改 `_process_store_event()`, 先尝试从 `_store_event_to_blocks` 弹出 transfer; 若为 None 则从 abandoned 字典弹出; 若仍为 None 则忽略 (陈旧事件)。之后调用新增的 `_release_transfer_refs()` 释放 CPU/GPU 块引用, 重置 CPU 块的哈希并将其归还到空闲池。
4. 更新连接器接口: 在 `SimpleCPUOffloadConnector.reset_cache()` 中将 `raise NotImplementedError` 替换为 `return self.scheduler_manager.reset()`, 并添加注释说明 worker 侧不主动接触, in-flight 传输会自然完成, 陈旧完成事件会被忽略。
5. 添加单元测试: 在 `test_scheduler.py` 中新增三个测试函数, 分别验证 eager 模式、lazy 模式和 pending load 场景下 `reset()` 的正确性。测试中构建 in-flight 传输, 调用 `reset()`

确认返回 `False`，模拟完成事件后检查块引用释放和前缀缓存可重置。

关键文件：

- `vllm/v1/simple_kv_offload/manager.py`（模块 KV 卸载；类别 `source`；类型 `core-logic`；符号 `_release_transfer_refs`, `reset`）：核心实现，包含 `reset` 逻辑、`abandoned` 状态追踪和引用释放
- `tests/v1/simple_kv_offload/test_scheduler.py`（模块 测试；类别 `test`；类型 `test-coverage`；符号 `test_reset_pending_eager_stores`, `test_reset_pending_lazy_stores`, `test_reset_pending_loads`）：新增三个测试用例，验证 `reset` 在各种模式下的正确性
- `vllm/distributed/kv_transfer/kv_connector/v1/simple_cpu_offload_connector.py`（模块 连接器；类别 `source`；类型 `core-logic`）：将 `reset_cache` 从 `NotImplementedError` 改为委托实现

关键符号：`reset`, `_release_transfer_refs`, `_process_store_event`, `has_pending_stores`, `reset_cache`

关键源码片段

`vllm/v1/simple_kv_offload/manager.py`

核心实现，包含 `reset` 逻辑、`abandoned` 状态追踪和引用释放

```
def reset(self) -> bool:
    """Abort all pending transfers, release block refs, reset CPU cache.
    Returns False if there are in-flight DMA transfers that must complete first.
    """
    gpu_pool = self._gpu_block_pool
    assert gpu_pool is not None, "reset() called before GPU pool bound"

    # 1. Move all pending store events to abandoned set
    # (their refs remain pinned until worker confirms completion)
    self._abandoned_store_event_to_blocks.update(
        self._store_event_to_blocks
    )
    self._store_event_to_blocks.clear()

    # 2. Move all pending loads to abandoned set
    self._abandoned_reqs_to_load.update(self._reqs_to_load)
    self._reqs_to_load.clear()
    self._load_event_to_reqs.clear()

    # 3. Clear scheduler-side queues and reset CPU prefix cache
    self._pending_cpu_hits.clear()
    self.cpu_block_pool.reset_prefix_cache()

    # 4. In lazy mode, reset the cursor
    if self._lazy_mode:
        self._cursor = None
```

```

# 5. Return True iff no in-flight transfers remain
return not (self._abandoned_store_event_to_blocks
            or self._abandoned_reqs_to_load)

def _release_transfer_refs(self, transfer: TransferMeta) -> None:
    """Release GPU/CPU block refs after a transfer (store or load) completes.
    CPU block hashes are cleared so the data won't be treated as valid cache.
    """
    cpu_blocks = [self.cpu_block_pool.blocks[bid] for bid in transfer.cpu_block_ids]
    for cb in cpu_blocks:
        cb.reset_hash()
    self.cpu_block_pool.free_blocks(cpu_blocks)

    assert self._gpu_block_pool is not None
    self._gpu_block_pool.free_blocks(
        [self._gpu_block_pool.blocks[bid] for bid in transfer.gpu_block_ids]
    )

```

tests/v1/simple_kv_offload/test_scheduler.py

新增三个测试用例，验证 reset 在各种模式下的正确性

```

# -----
# Test 12: Reset with pending eager stores waits for completion
# -----
def test_reset_pending_eager_stores() -> None:
    """Eager mode: reset() abandons in-flight stores until they complete."""
    fix = make_scheduler(num_cpu_blocks=8, num_gpu_blocks=16, lazy=False)
    sched = fix.scheduler
    gpu_pool = fix.gpu_block_pool

    num_blocks = 2
    req = make_request(num_blocks=num_blocks)

    kv_blocks = _alloc_and_register(fix, req, num_blocks)
    sched.update_state_after_alloc(req, kv_blocks, num_external_tokens=0)
    block_ids = kv_blocks.get_block_ids()
    sched_out = make_scheduler_output(
        {req.request_id: num_blocks * BLOCK_SIZE},
        new_reqs={req.request_id: block_ids},
    )

    meta = sched.build_connector_meta(sched_out)
    assert meta.store_event >= 0
    assert len(sched._store_event_to_blocks) > 0

    # GPU blocks should have elevated ref_cnt from touch()
    for bid in meta.store_gpu_blocks:
        assert gpu_pool.blocks[bid].ref_cnt > 0

    # Free the request's own block refs (simulates preemption)

```

```

gpu_pool.free_blocks(gpu_pool.blocks[bid] for bid in block_ids[0])

# Reset should keep DMA refs pinned until the worker reports completion.
assert sched.reset() is False
assert len(sched._store_event_to_blocks) == 0
assert len(sched._abandoned_store_event_to_blocks) == 1
assert len(sched._reqs_to_store) == 0
assert len(sched._store_event_to_reqs) == 0

num_used = gpu_pool.num_gpu_blocks - gpu_pool.get_num_free_blocks()
assert num_used > 1

simulate_store_completion(sched, meta.store_event)
assert len(sched._abandoned_store_event_to_blocks) == 0

# All GPU blocks should now be free (ref_cnt == 0) except null block
num_used = gpu_pool.num_gpu_blocks - gpu_pool.get_num_free_blocks()
assert num_used == 1, f"Expected only null block in use, got {num_used}"

# GPU prefix cache reset should now succeed
assert gpu_pool.reset_prefix_cache() is True
assert sched.reset() is True

```

[vllm/distributed/kv_transfer/kv_connector/v1/simple_cpu_offload_connector.py](#)

将 `reset_cache` 从 `NotImplementedError` 改为委托实现

```

# NOTE: Workers are not contacted. In-flight transfers drain naturally,
# and stale completions are ignored by the guarded
# SimpleCPUOffloadScheduler._process_store_event().
def reset_cache(self) -> bool | None:
    if self.scheduler_manager is not None:
        return self.scheduler_manager.reset()
    return None

```

评论区精华

- aoshen02 对同步安全性的质疑：' 仅清除调度器侧状态不足，worker 侧异步传输可能仍在运行。' jonathanc-n 通过引入 `abandoned` 字典保证引用不释放，`reset` 在传输完成前返回 `False`。方案被接受。
- gemini-code-assist[bot] 关于重复 `free` 的建议：使用 `set` 收集 `block ID` 防止重复 `free`。jonathanc-n 认为当前逻辑已保证去重，ivanium 指出存在跨步重复 bug，将在 PR#41289 中修复，并建议在 `reset` 中清理 `_in_flight_store_gpu_blocks`。该建议被部分采纳，等待 `dedup` 修复合并。
- NickLucche 对保护条件的建议：将 `gpu_pool is None` 的静默返回改为断言，因为 `reset` 不应在调度器未初始化时调用。作者随后改为断言。

- 确保 worker 侧异步传输完成后再释放资源 (design): 采用 abandoned 方案, 等待传输自然完成, 不主动同步 worker。
- 使用 set 避免 reset 中重复 free 导致引用计数错误 (correctness): 待合并 dedup 修复, 作者可能需要在 reset 中补充清理 `_in_flight_store_gpu_blocks`。
- 将 `gpu_pool is None` 检查改为断言 (design): 改为断言, 更早暴露状态错误。
- worker 是否需要同步 reset (design): worker 侧不处理, 依赖自然完成 + abandoned 机制。

风险与影响

• 风险:

1. 依赖外部修复: reset 的正确性依赖于即将合并的 dedup 修复 (PR#41289), 若未合并或修复不完整, 可能导致引用计数异常。
2. 异步竞争: 虽然 abandoned 机制防止了传输中释放, 但未同步 worker, 如果 worker 在 reset 后立即发送完成事件, abandoned 字典可能已被清除, 但事件被正确处理 (忽略)。但若多个 reset 连续调用, 可能造成状态混乱。
3. 测试覆盖不足: 测试仅覆盖单次 reset, 缺少多级 reset 嵌套或并发场景。
4. 资源泄漏: 如果传输永远不完成 (异常情况), abandoned 条目会永久存在, 可能导致内存泄漏。但生产环境中传输超时会保障有限等待。 - 影响: 影响范围限定在使用 SimpleCPUOffloadConnector 的模型推理路径。用户现在可以调用 `reset_cache()` 安全重置 CPU 前缀缓存, 等待所有 in-flight DMA 传输完成后才真正释放 GPU/CPU 块。对其他连接器无影响, API 兼容性未破坏 (之前是 `NotImplementedError`, 现在是正常返回)。 - 风险标记: 依赖 dedup 修复, 异步传输竞争, 测试覆盖有限

关联脉络

- PR #41289 Fix dedup of GPU blocks in eager offloading: ivanium 在 review 中指出该 PR 修复了跨步重复 bug, 要求在 reset 中清理 `_in_flight_store_gpu_blocks` 以确保正确性。