

PR #29843 完整报告

sgl-project/sglang

[trtllm_mha] Fuse cuda-graph metadata rebuild into one triton kernel

合并时间: 2026-07-04 13:24

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/29843>

执行摘要

本 PR 将 TRTLLM MHA 后端的 CUDA graph 元数据重建路径从约 25 个分散的 ATen 操作融合为一个 Triton 内核，显著降低 host CPU 调度开销和 TP 秩间 jitter。同时修复了混合注意力后端和推测解码运行器中的集成缺失，并添加了全面的测试套件。该变更仅影响 Blackwell 平台。

功能与动机

原有实现在每次 graph replay 时，使用约 25 个小 ATen 操作重建 cache_seqlens、cu_seqlens_k/q、page_table 等元数据。这些操作产生大量 host dispatch，在某些 CPU 上需 0.7-1.0 ms，且在 draft-decode、target-verify、draft-extend 多个阶段重复，导致各 TP 秩的 cudaGraphLaunch 时间偏差，该偏差被计入第一个 all-reduce 的等待时间。本 PR 旨在消除此 CPU jitter，提升多秩协同效率。

实现拆解

1. 新增融合 Triton 内核：在 trtllm_mha_graph_metadata.py 中定义 update_trtllm_mha_graph_metadata_kernel，利用每个 batch 行一个 program 的并行模式，一次计算所有元数据。支持三种 q_mode（预设 /cumsum/strided）、SWA 翻译和 -1 sentinel 保护。
2. 后端集成替换：在 trtllm_mha_backend.py 中导入新内核，重写 _apply_cuda_graph_metadata，将原先约 25 个操作替换为单次内核调用。同时预处理 SWA 映射 (_swa_full_to_swa_mapping) 供内核使用。init_forward_metadata_in_graph 现在调用融合内核作为 graph 录制的一部分。
3. 推测解码运行器扩展：在 eagle_draft_extend_cuda_graph_runner.py、frozen_kv_mtp_cuda_graph_runner.py、multi_layer_eagle_draft_extend_cuda_graph_runner.py 的 capture_one_shape 中插入 init_forward_metadata_in_graph 调用。draft-extend 路径使用静态捕获的 q_stride 避免运行时 .item() 调用。
4. 混合注意力后端修复：hybrid_attn_backend.py 和 hybrid_linear_attn_backend.py 原未转发 in-graph 钩子，本 PR 添加 init_forward_metadata_in_graph 方法以委托给内部后端。
5. 测试配套：新增 test_trtllm_mha_graph_metadata.py (432 行)，包含基于纯 ATen 参考的正确性测试（覆盖 batch 组合、q_mode、SWA 开关、零填充尾部），以及钩子分发测试和 graph 录制重放测试。现有 test_trtllm_mha.py 增加 frozen_kv_mtp 的 runner 测试。

以下片段来自测试文件 `test_trtllm_mha_graph_metadata.py`，展示如何验证融合内核通过钩子被图录制调用：

```
def test_cuda_graph_metadata_launch_runs_in_graph_hook(monkeypatch):
    calls = []
    def fake_update(**kwargs):
        calls.append(kwargs)
        monkeypatch.setattr(
            trtllm_mha_backend, "update_trtllm_mha_graph_metadata", fake_update
        )
    backend = _make_backend_for_hook_test()
    fb = SimpleNamespace(
        batch_size=2,
        req_pool_indices=torch.arange(2, dtype=torch.int64),
        seq_lens=torch.ones(2, dtype=torch.int32),
        forward_mode=ForwardMode.DECODE,
        spec_info=None,
        positions=torch.arange(2, dtype=torch.int64),
        out_cache_loc=torch.arange(2, dtype=torch.int64),
    )
    backend.init_forward_metadata_out_graph(fb,
        in_capture=True)
    assert calls == []
    assert backend.forward_metadata is backend.decode_cuda_graph_metadata[2]
    backend.init_forward_metadata_in_graph(fb)
    assert len(calls) == 1
    assert calls[0]["out_cache_loc"] is fb.out_cache_loc
    calls.clear()
    backend.init_forward_metadata_out_graph(fb)
    assert calls == []
    assert backend.forward_metadata is backend.decode_cuda_graph_metadata[2]
```

该测试使用 monkeypatch 拦截内核调用，验证 `init_forward_metadata_out_graph`（graph 外部准备）不触发内核，而 `init_forward_metadata_in_graph`（graph 内部）仅触发一次。这确保了融合内核被正确地录制到 CUDA graph 中而非在 capture 前执行。

评论区精华

- Qiaolin-Yu指出 hybrid 后端缺少 `init_forward_metadata_in_graph` 转发，导致融合内核在混合注意力模型中不生效。pranjalssh随即添加了对应转发，解决了此问题。
- Qiaolin-Yu询问 `frozen_kv_mtp` 是否已经测试。pranjalssh添加了专用的 runner 测试用例。
- Qiaolin-Yu就 `multi_layer_eagle` 的测试提出同样问题，pranjalssh询问推荐测试方法，目前尚无明确解决方案。
- ch-wan质疑 SWA 支持是否在融合后丢失。merrymercy解答：SWA 翻译已通过 `_apply_cuda_graph_metadata` 中的融合内核处理，无需额外操作。

风险与影响

风险

- SWA 映射一致性：SWA 依赖在构造时捕获的 `full_to_swa_mapping`，若运行时池变化未更新，可能导致翻译错误，但 SWA 池通常稳定。
- `multi_layer_eagle` 测试空白：该路径无直接测试覆盖，存在回归风险。现有 EAGLE 链式测试提供部分保障。
- Blackwell 独占：新内核仅 sm100 可用，其他 GPU 仍使用原有 ATen 路径，功能正常但无优化。
- 混合后端转发完整性：若新增后端类型未正确转发，in-graph 钩子可能静默缺失，但现有测试覆盖了 hybrid 后端的基本路径。

影响

- 用户：Blackwell 用户在多 TP 场景下应感受到更低的延迟抖动和可能的吞吐提升。其他 GPU 用户无影响。
- 系统：减少 host-device 交互频率，降低 CPU 占用，使 cudaGraphLaunch 跨秩更同步。
- 团队：需维护新的 Triton kernel，并确保与未来后端演进兼容。

关联脉络

本 PR 与近期 #30025 (DSA indexer 双流修复) 和 #30088 (禁用 DSA indexer fusion) 均属于注意力后端 CUDA graph 优化系列，反映团队在减少 graph 录制 / 重放开销方面的持续投入。TRTLLM MHA 与 DSA 后端采用不同的融合策略，但目标一致：消除 host 侧的元数据重建瓶颈。