

PR #27233 完整报告

sgl-project/sglang

[Spec] Fuse small kenrels under `gather\_spec\_extras`

合并时间: 2026-06-09 06:02

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/27233>

执行摘要

- 一句话: 融合 Spec 中 4 个 gather 算子, TPOT 提升 1-2%
- 推荐动作: 值得精读, 展示了如何用 Triton 手动融合 gather 操作绕过 torch.compile 融合限制, 并附有高质量的测试设计, 可作为 Triton kernel 开发的范例。

功能与动机

Torch compile 无法自动融合 4 个独立的 gather 操作, 导致在 CPU 慢的机器上编译图延迟高达 100us。手动编写 Triton kernel 进行融合以降低延迟, 提升推理速度。

实现拆解

1. 在 python/sglang/srt/speculative/triton_ops/gather_spec_extras.py 中创建 Triton kernel `_gather_rows_kernel`, 通过一次 kernel launch 按行 gather 所有 buffer。
2. 编写封装函数 `gather_spec_extras`, 处理参数预处理、输出分配和 kernel 调用。
3. 在 `overlap_utils.py` 中删除旧的 `_gather_spec_extras` 函数 (基于 `@torch.compile`), 导入新的 `gather_spec_extras` 并替换 `FutureMap._resolve_spec_extras` 中的调用。
4. 新增测试文件 `test_gather_spec_extras.py`, 包含参考实现 `_ref_gather` 和全面测试用例, 覆盖多种 shape、空索引、非连续索引等边界条件, 并验证数值比特一致性及源 buffer 无写副作用。

关键文件:

- `python/sglang/srt/speculative/triton_ops/gather_spec_extras.py` (模块 Triton 内核; 类别 source; 类型 core-logic; 符号 `_gather_rows_kernel`, `_row_width`, `_empty_like_rows`, `gather_spec_extras`): 新增 Triton 融合 gather 内核, 通过一次 kernel launch 完成 `topk_p`、`topk_index`、`bonus_tokens` 及可选 `hidden_states` 的 gather, 是本次性能优化的核心。
- `test/registered/kernels/test_gather_spec_extras.py` (模块 单元测试; 类别 test; 类型 test-coverage; 符号 `_ref_gather`, `_make_buffers`, `TestGatherSpecExtras`, `setUpClass`): 为新 Triton kernel 提供全面单元测试, 包含参考实现、多种 shape 组合、空索引、非连续索引等边界条件, 确保数值一致性和无副作用。
- `python/sglang/srt/managers/overlap_utils.py` (模块 调度器; 类别 source; 类型 core-logic; 符号 `_gather_spec_extras`, `gather_spec_extras`): 删除旧的 `torch.compile` 版本, 导入并使用新的 Triton kernel, 是集成变更的关键文件。

关键符号: gather_spec_extras, _gather_rows_kernel, _ref_gather, _assert_matches_reference

关键源码片段

[python/sglang/srt/speculative/triton_ops/gather_spec_extras.py](#)

新增 Triton 融合 gather 内核, 通过一次 kernel launch 完成 topk_p、topk_index、bonus_tokens 及可选 hidden_states 的 gather, 是本次性能优化的核心。

```
def gather_spec_extras(
    indices: torch.Tensor,
    topk_p_buf: torch.Tensor,
    topk_index_buf: torch.Tensor,
    output_tokens_buf: torch.Tensor,
    hidden_states_buf: Optional[torch.Tensor],
):
    # 确保索引连续 (若源自不同生产者可能非连续)
    indices = indices.contiguous()
    m = indices.shape[0]
    has_hidden = hidden_states_buf is not None

    # 预分配输出缓冲区 (形状来自源 buffer 但行数改为 m)
    topk_p = _empty_like_rows(topk_p_buf, m)
    topk_index = _empty_like_rows(topk_index_buf, m)
    bonus_tokens = _empty_like_rows(output_tokens_buf, m)
    hidden_states = _empty_like_rows(hidden_states_buf, m) if has_hidden else None
    if m == 0:
        return topk_p, topk_index, bonus_tokens, hidden_states

    # 每行的元素个数 (展平后), 1-D buffer 的 row_width = 1
    n0 = _row_width(topk_p_buf)
    n1 = _row_width(topk_index_buf)
    n2 = _row_width(output_tokens_buf)
    n3 = _row_width(hidden_states_buf) if has_hidden else 1
    max_n = max(n0, n1, n2, n3)

    # 计算列块数 (每个块 BLOCK 列)
    BLOCK = 1024
    grid = (m, triton.cdiv(max_n, BLOCK))
    # 确保 buffer 指针传给 kernel 时即使无对应 buffer 也有有效占位 (dummy 指针)
    dummy = topk_p_buf.flatten().data_ptr() # 仅作占位
    _gather_rows_kernel[grid](
        indices,
        topk_p_buf, topk_p,
        topk_index_buf, topk_index,
        output_tokens_buf, bonus_tokens,
        hidden_states_buf if has_hidden else dummy,
        hidden_states if has_hidden else dummy,
        n0, n1, n2, n3,
```

```

        HAS3=has_hidden,
        BLOCK=BLOCK,
    )
    return topk_p, topk_index, bonus_tokens, hidden_states

```

test/registered/kernels/test_gather_spec_extras.py

为新 Triton kernel 提供全面单元测试，包含参考实现、多种 shape 组合、空索引、非连续索引等边界条件，确保数值一致性和无副作用。

```

_OUTPUT_NAMES = ("topk_p", "topk_index", "bonus_tokens", "hidden_states")

def _ref_gather(indices, topk_p_buf, topk_index_buf, output_tokens_buf, hidden_states_buf):
    # 参考实现：使用 torch 高级索引 gather，与被替换的 torch.compile 路径保持一致
    topk_p = topk_p_buf[indices]
    topk_index = topk_index_buf[indices]
    bonus_tokens = output_tokens_buf[indices]
    hidden_states = hidden_states_buf[indices] if hidden_states_buf is not None else None
    return topk_p, topk_index, bonus_tokens, hidden_states

class TestGatherSpecExtras(CustomTestCase):
    def _assert_matches_reference(self, indices, bufs):
        # 克隆源 buffer 以便后续检查是否有副作用
        src_snapshots = [None if b is None else b.clone() for b in bufs]
        ref = _ref_gather(indices, *bufs)
        got = gather_spec_extras(indices, *bufs)
        self.assertEqual(len(got), len(ref))
        for name, r, o in zip(_OUTPUT_NAMES, ref, got):
            if r is None:
                self.assertIsNone(o, f"{name} should be None when buffer is None")
                continue
            # 纯 gather 应为比特精确拷贝
            torch.testing.assert_close(o, r, rtol=0, atol=0, msg=f"{name} value mismatch")
        # 确保源 buffer 未被修改
        for name, before, buf in zip(_OUTPUT_NAMES, src_snapshots, bufs):
            if before is None:
                continue
            torch.testing.assert_close(buf, before, rtol=0, atol=0, msg=f"source buffer {name} was mutated")

```

python/slang/srt/managers/overlap_utils.py

删除旧的 torch.compile 版本，导入并使用新的 Triton kernel，是集成变更的关键文件。

```

def _resolve_spec_extras(self, batch: ScheduleBatch) -> None:
    draft_input: EagleDraftInput = batch.spec_info
    if draft_input is None:
        return
    indices = draft_input.future_indices
    indices.record_stream(torch.get_device_module(self.device).current_stream())
    hidden_states_buf = self.hidden_states_buf if spec_need_hidden_states() else None

```

```
# 使用新 Triton kernel 替换原来的 torch.compile 版本
draft_input.topk_p, draft_input.topk_index, draft_input.bonus_tokens, hidden_states = gather_spec_extras(
    indices,
    self.topk_p_buf,
    self.topk_index_buf,
    self.output_tokens_buf,
    hidden_states_buf,
)
# 后续处理 hidden_states...
```

评论区精华

审核者 Qiaolin-Yu 要求增加单元测试以确保正确性，作者随后添加了覆盖多种 shape 和边界条件的测试用例，得到审核者认可并批准。

- 缺少单元测试 (testing): 作者增加了 test_gather_spec_extras.py，包含参考实现对比和多种边界情况测试。审核者认可并批准。

风险与影响

- 风险：该变更用手写 Triton kernel 替换 torch.compile 版本，数值一致性风险较低，因为测试用例严格对比输出（零容差），并验证源 buffer 未被修改。潜在风险包括：新 kernel 依赖 Triton（但已是 sglang 默认依赖）；假设源 buffer 为 row-contiguous（分配时保证，但若未来修改 layout 可能失效）；边界情况如空索引、非连续索引已在测试中覆盖。总体回归风险较小。
- 影响：对用户：TPOT 小幅提升（1-2%），无功能变化。对系统：减少 kernel launch 数量，降低 CPU 开销。对团队：新增一个 Triton kernel 文件和一个测试文件，维护成本低。
- 风险标记：新 Triton 内核数值一致性，边界条件（空索引 / 非连续索引），假设源 buffer 连续，Triton 依赖

关联脉络

- PR #27599 [Spec] Naming cleanup: contiguous draft-loc kernel + accepted->accept: 同一 speculative-decoding 功能线，共享 overlap_utils 及 triton_ops 模块，此 PR 为推理性能优化延续。
- PR #27552 [Spec] Rename token resolver to _resolve_spec_v2_tokens; remove dead V1 helpers: 同样涉及 speculation 路径的清理与重构，与本次融合 kernel 的上下文紧密相关。