

PR #35059 完整报告

sgl-project/sglang

[Spec] Resolve shared-read ends from the backend declaration alone

合并时间: 2026-08-17 16:35

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/35059>

执行摘要

- 一句话: 共享读终点改为仅按 backend 声明解析, 删除冗余 phase gate
- 推荐动作: 值得精读。该 PR 是典型的设计简化: 通过识别两个冗余 gate, 将复杂的哪个阶段发布共享读完成事件收敛到 `last_shared_read_runner` 单一事实源; 同时展示了算法特异性约束 (NGRAM) 如何迫使保留 coarse fence。对于理解 SGLang 的 WAR barrier、speculative decoding 与 CUDA graph runner 的协作有直接帮助。关注点: 删除公共方法的迁移影响、NGRAM 清除逻辑的长期维护。

功能与动机

PR body 指出 decode runner 在询问 backend 之前施加的双重 gate 均已冗余: 一是 `elif not forward_mode.is_decode()`, 因为到达该 runner 的模式只有 decode/target verify 等, 且基础 `shared_read_ends` 已对非 decode/target-verify 返回 `UNKNOWN`; 二是 `is_last_shared_read_phase`, 因为哪个 runner 拥有 publish 已由 `last_shared_read_runner` 回答, 该方法没有任何其他调用者。对于 NGRAM 这种 scheduler-read runner 就是 target runner 的算法, 去掉 gate 会把不安全的 fine-grained verify 记录交给 barrier, 复现 `test_spec_ngram.py` 的设备端索引断言, 因此 NGRAM 需要特殊处理。整个变更的动机是消除冗余控制流, 并让 read-end 归属只由 backend 声明和 `last_shared_read_runner` 决定。

实现拆解

1. 移除 decode runner 的 phase gate: 在 `python/sglang/srt/model_executor/runner/decode_cuda_graph_runner.py` 的 `_resolve_shared_read_ends()` 中删除 `forward_mode.is_target_verify()`、`is_last_shared_read_phase()` 检查与 `not forward_mode.is_decode()` 分支, 改为直接调用 `attn_backend.shared_read_ends(forward_mode)`; 仅保留 `IN_REPLAY` 且无 `in_graph_metadata_prep_done` 时降级为 `PRE_REPLAY` 的兜底逻辑。原因: base 实现已对非 decode/target-verify 返回 `UNKNOWN`, 且发布者已由 `last_shared_read_runner` 唯一决定。
2. 删除无调用者的 predicate: 从 `SpeculativeAlgorithm (spec_info.py)` 与 `CustomSpecAlgo (spec_registry.py)` 中删除 `is_last_shared_read_phase()`。该方法唯一的调用点在上一步被删除后即无引用。
3. NGRAM 特殊处理: 在 `python/sglang/srt/speculative/ngram_worker.py` 的 `forward_batch_generation()` verify 返回后, 将 `self.target_worker.model_runner.share`

`d_read_done_event` 置为 `None`，使调度器对 NGRAM 继续使用 `coarse whole-forwardfence`，避免设备端 `indexSelectSmallIndex` 断言。

4. 注释与文档同步：`scheduler.py` 的 `_apply_war_barrier()` 与 `model_runner.py` 的 `shared_read_done_event` 字段注释更新为 `read-done mailbox` 语义，明确事件由 `last_shared_read_runner` 提供、`None` 表示 `coarse fence`。
5. 测试配套：`test_decode_cuda_graph_shared_read_fence.py` 移除 `mode/owns_verify` 参数维度，参数化精简为 `declared/has_marker`，并新增断言确保 `_resolve_shared_read_ends` 只调用一次 `shared_read_ends(Decode)`，将新契约固化为回归测试。

关键文件：

- `python/sglang/srt/model_executor/runner/decode_cuda_graph_runner.py`（模块 图执行器；类别 `source`；类型 `data-contract`；符号 `_resolve_shared_read_ends`）：核心变更点：删除 `_resolve_shared_read_ends` 中的 `forward mode gate` 与 `spec_algorithm` 检查，改为直接采用 `backend` 声明。
- `python/sglang/srt/speculative/ngram_worker.py`（模块 Ngram 执行器；类别 `source`；类型 `core-logic`；符号 `forward_batch_generation`）：NGRAM 特殊处理：`verify` 后清空 `read-done` 事件，避免设备端索引断言，是 `gate` 删除后唯一需要定向保护的算法。
- `test/registered/unit/model_executor/runner/test_decode_cuda_graph_shared_read_fence.py`（模块 共享读栅栏；类别 `test`；类型 `test-coverage`；符号 `_runner`，`test_resolve_shared_read_ends`）：测试核心行为，参数化精简并断言只调用一次 `backend.shared_read_ends`，固化新契约。
- `python/sglang/srt/speculative/spec_info.py`（模块 投机算法；类别 `source`；类型 `core-logic`；符号 `is_last_shared_read_phase`）：从 `SpeculativeAlgorithm` 中删除无调用方的 `is_last_shared_read_phase` 方法，是 API 精简的一部分。
- `python/sglang/srt/speculative/spec_registry.py`（模块 算法注册；类别 `source`；类型 `core-logic`；符号 `is_last_shared_read_phase`）：同步删除 `CustomSpecAlgo` 中的 `is_last_shared_read_phase`，保持基类与插件基类一致。
- `python/sglang/srt/managers/scheduler.py`（模块 调度器；类别 `source`；类型 `documentation`）：更新 `_apply_war_barrier` 注释，明确事件由 `last_shared_read_runner` 提供，而非固定阶段列表。
- `python/sglang/srt/model_executor/model_runner.py`（模块 模型执行器；类别 `source`；类型 `documentation`）：更新 `shared_read_done_event` 字段注释，反映 `read-done mailbox` 语义，降低后续维护困惑。

关键符号：`_resolve_shared_read_ends`，`forward_batch_generation`，`_apply_war_barrier`，`is_last_shared_read_phase`，`test_resolve_shared_read_ends`

关键源码片段

`python/sglang/srt/model_executor/runner/decode_cuda_graph_runner.py`

核心变更点：删除 `_resolve_shared_read_ends` 中的 `forward mode gate` 与 `spec_algorithm` 检查，改为直接采用 `backend` 声明。

```
def _resolve_shared_read_ends(self, attn_backend, forward_mode) -> SharedReadEnds:
    # 信任 attention backend 的声明：只有它知道共享缓冲区的读终点在哪里。
    declared = attn_backend.shared_read_ends(forward_mode)

    if (
        declared is SharedReadEnds.IN_REPLAY
        and self.in_graph_metadata_prep_done is None
    ):
        # TODO: PRE_REPLAY 比声明更早落地，POST_REPLAY 才是更安全的选择。
        # 当前没有 in-graph marker 可记录事件，只能退回 pre-replay 粗粒度记录。
        return SharedReadEnds.PRE_REPLAY
    return declared
```

python/sclang/srt/speculative/ngram_worker.py

NGRAM 特殊处理：verify 后清空 read-done 事件，避免设备端索引断言，是 gate 删除后唯一需要定向保护的算法。

```
# NGRAMWorker.forward_batch_generation 的 verify 分支
batch_result = self.target_worker.forward_batch_generation(
    batch, is_verify=True
)
# NGRAM 的 verify 会在 runner 的 in-graph marker 之后继续读取
# scheduler 共享状态，因此不能发布 fine-grained read-done 事件。
# 显式清空事件，让调度器回退到 coarse whole-forward fence，
# 避免设备端 indexSelectSmallIndex 断言。
self.target_worker.model_runner.shared_read_done_event = None
```

test/registered/unit/model_executor/runner/test_decode_cuda_graph_shared_read_fence.py

测试核心行为，参数化精简并断言只调用一次 backend.shared_read_ends，固化新契约。

```
def test_resolve_shared_read_ends(declared, has_marker, expected):
    runner = _runner(has_marker=has_marker)
    backend = _backend(declared)

    # 新契约：resolve 只看 backend 声明与是否有 in-graph marker。
    assert runner._resolve_shared_read_ends(backend, DECODE) is expected
    # 且只查询一次 backend，不依赖额外的 forward-mode 门控。
    backend.shared_read_ends.assert_called_once_with(DECODE)
```

评论区精华

该 PR 没有 review 评论。合并前的 CI 反馈线程：作者执行 `/rerun-test test_spec_ngram.py test_spec_eagle.py test_dflash.py test_mimo_v2_flash.py test_decode_cuda_graph_shared_read_fence.py`，机器人返回 1-gpu-h100 上 `test_spec_ngram.py` 失败，其余 4 项通过。这与 PR body 的分析一致——NGRAM 是 gate 唯一承重的算法；作者随后在 commit `fix ngram` 中通过清除事件恢复 coarse fence。体现了测试先行暴露、随后定向修复的闭环。

- test_spec_ngram.py 在 1-gpu-h100 上失败及修复 (testing): 作者在 commit 'fix ngram' 中让 NGRAMWorker 在 verify 后清空 shared_read_done_event, 恢复 coarse fence, 测试相应恢复。

风险与影响

- 风险:

1. 插件兼容性 (breaking change) : is_last_shared_read_phase 是 SpeculativeAlgorithm/CustomSpecAlgo 的公共方法, 外部自定义算法若覆盖或调用它, 需要迁移到 last_shared_read_runner 语义, 否则升级后可能 AttributeError 或行为变化。
2. 自定义 attention backend 暴露: 删除 gate 后, decode runner 会对所有到达的 forward mode 调用 shared_read_ends()。若某个后端在非 decode/target-verify 模式返回非 UNKNOWN 声明, 可能产生意外的 fine-grained fence。PR 断言 in-tree 无 override 偏离, 但外部插件需要自查。
3. NGRAM 清除逻辑缺乏守护: ngram_worker.py 的 = None 是一个隐式约定, 若未来有人删除该行, 会重新触发设备端索引断言, 且没有显式测试在 CI 中守护这一约定 (test_spec_ngram.py 已在 CI 中)。
4. 影响面: 变更位于 decode cuda graph runner 与调度器 WAR barrier 之间, 属于 speculative decoding 核心路径, 回归风险集中在 fence 粒度错误导致的悬空读写。
 - 影响: 用户与系统层面: speculative decoding 的共享读 fence 粒度更精确, 多数算法 (如 EAGLE 系列) 的事件发布由 last_shared_read_runner 唯一决定, 减少了对算法枚举的依赖; NGRAM 保持 coarse fence, 行为不变且规避设备端断言。团队层面: 删除了一个公共 API, 需要同步插件文档; 注释更新降低了后续维护者对谁发布事件的困惑。影响范围为 speculative 解码 + CUDA graph + 调度器, 通常由 CI 中的 test_spec_* 与 test_decode_cuda_graph_shared_read_fence 覆盖。
 - 风险标记: 删除公共 API (is_last_shared_read_phase), NGRAM 依赖手动清除事件, 核心调度路径变更, 外部插件 /backend 兼容性

关联脉络

- PR #35057 Unknown (stacks-on PR mentioned in body): PR body 明确说明本 PR stacks on #35057, 后者使 verify record 对 multi-layer eagle 无害, 是删除 phase gate 的前提。