

PR #31488 完整报告

sgl-project/sglang

Overlap grammar (constrained decoding) with speculative decode verify

合并时间: 2026-07-21 08:36

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/31488>

执行摘要

- 一句话: Grammar 约束解码与推测解码 verify 阶段重叠, 消除 GPU 空闲气泡
- 推荐动作: 值得精读。设计思路 (将 CPU 工作移到 GPU forward 窗口内) 和实现细节 (异步 D2H + CUDA event + pin-memory barrier) 是高性能推理系统的典型技巧。代码改动清晰, 注释详尽, 适合作为性能优化的参考案例。

功能与动机

Speculative decoding + grammar 时, 每个 verify 步骤需要 CPU 往返: 拷贝 draft tokens (D2H)、推进 grammar FSM、构建 vocab bitmask 并拷贝回 GPU (H2D)。这些操作目前是阻塞的, 且调度器 `need_grammar_sync` 强制关闭跨 batch 重叠, 导致 GPU 在每个 grammar decode 步骤都有大段空闲。PR 的目标是将 grammar CPU 工作 (FSM 推进和 bitmask 构建) 移到 GPU verify forward 的窗口内, 从而消除空闲气泡。

实现拆解

1. 添加能力标记: 在 `SpeculativeAlgorithm` 类上添加 `supports_grammar_overlap()` 方法 (EAGLE 返回 True, CustomSpecAlgo 默认 False)。调度器根据此标记决定是否强制关闭重叠。
2. 调度器保持重叠: 在 `scheduler.py` 的 `is_disable_overlap_for_batch` 中, 当 spec 算法支持 grammar overlap 时, 不再因为 grammar 而关闭重叠。新增 `_advance_pending_grammar()` 方法, 它遍历 `result_queue` 中尚未处理的 batch, 调用 `batch_result_processor.advance_grammar_fsm()` 推进 FSM。该方法作为 `grammar_barrier` 传递给 worker。
3. 统一 FSM 推进点: 在 `batch_result_processor.py` 中引入 `advance_grammar_fsm()` 方法, 处理 decode 和 extend 的 FSM 推进。decode 时使用已记忆的 `grammar_retained_tokens` 避免重复推进。`_apply_prefill_grammar` 增加 `already_advanced` 参数以支持已被 barrier 推进的情况。
4. Eagle worker 重叠拷贝与 forward: 在 `eagle_worker_common.py` 的 `run_eagle_verify` 中, 将 draft/verify D2H 拷贝改为异步 (使用 pin-memory 缓冲区 + CUDA event), 在 target verify forward 发射后、bitmask 构建前同步 event。支持可选的 `grammar_barrier` 回调, 在 target forward 之后、bitmask 之前调用。bitmask 的 H2D 拷贝使用 `non_blocking=True`。

5. XGrammar pinned bitmask: 在 `xgrammar_backend.py` 中新增 `_allocate_token_bitmask()` 函数, 分配 pin-memory 的 bitmask 张量, 使后续 H2D 拷贝真正异步。
6. 配套改动: 更新 `GenerationBatchResult` 记录 `grammar_advanced` 和 `grammar_retained_tokens`; `ScheduleBatch.copy()` 传递 `has_grammar`; EAGLE 单层 / 多层 worker 和 FrozenKV MTP worker 的 `forward_batch_generation` 转发 `grammar_barrier`。测试文件补充了 fake 的 `forward_mode` 支持。

关键文件:

- `python/sglang/srt/managers/scheduler_components/batch_result_processor.py` (模块 批处理结果处理器; 类别 source; 类型 core-logic; 符号 `_apply_prefill_grammar`, `advance_grammar_fsm`): 核心实现: 新增 `advance_grammar_fsm` 统一 FSM 推进, `_apply_prefill_grammar` 支持 `already_advanced` 标记, 避免重叠路径下重复推进。同时 `_resolve_spec_v2_tokens` 中调用 `advance_grammar_fsm` 并利用 `grammar_retained_tokens` 避免二次 FSM 遍历。
- `python/sglang/srt/managers/scheduler.py` (模块 调度器; 类别 source; 类型 core-logic; 符号 `_advance_pending_grammar`): 调度器核心: 修改 `is_disable_overlap_for_batch`, 使支持 grammar overlap 的算法不再因 grammar 强制关闭重叠。新增 `_advance_pending_grammar` 作为 grammar barrier, 在 worker verify 之前推进上一个 batch 的 FSM, 实现 CPU 与 GPU 重叠。
- `python/sglang/srt/speculative/eagle_worker_common.py` (模块 Eagle Worker; 类别 source; 类型 dependency-wiring): 共享 verify 函数 `run_eagle_verify` 添加 `grammar_barrier` 参数, 实现异步 D2H 拷贝 + CUDA event, 使 grammar CPU 工作与 target verify forward 重叠。同时将 vocab mask 的 H2D 拷贝改为 `non_blocking`, 进一步隐藏延迟。
- `python/sglang/srt/constrained/xgrammar_backend.py` (模块 约束后端; 类别 source; 类型 core-logic; 符号 `_allocate_token_bitmask`): 为确保 `non_blocking` H2D 真正异步, `allocate_vocab_mask` 改为分配 pin-memory 张量, 替代默认 pageable 分配。这是实现重叠的关键基础设施。
- `python/sglang/srt/speculative/eagle_worker_v2.py` (模块 Eagle Worker; 类别 source; 类型 core-logic; 符号 `forward_batch_generation`, `verify`): 单层 EAGLE worker 入口 `forward_batch_generation` 转发 `grammar_barrier` 给 `verify` 方法, 进而传给 `run_eagle_verify`。
- `python/sglang/srt/speculative/multi_layer_eagle_worker_v2.py` (模块 Eagle Worker; 类别 source; 类型 core-logic; 符号 `forward_batch_generation`, `verify`): 多层 EAGLE worker 同步修改, 与单层对称。
- `python/sglang/srt/speculative/spec_info.py` (模块 推测配置; 类别 source; 类型 core-logic; 符号 `supports_grammar_overlap`): `SpeculativeAlgorithm` 基类添加 `supports_grammar_overlap` 方法 (默认返回 False), EAGLE 子类覆盖为 True。
- `python/sglang/srt/speculative/spec_registry.py` (模块 推测注册; 类别 source; 类型 core-logic; 符号 `supports_grammar_overlap`): 注册层同步添加 `supports_grammar_overlap` 标记。

- python/sglang/srt/managers/utils.py (模块 工具函数; 类别 source; 类型 core-logic) : 新增 `_async_d2h` 工具函数, 用于在 pin-memory 张量上执行非阻塞 device-to-host 拷贝。
- test/registered/unit/managers/test_batch_result_processor_spec_grammar.py (模块 单元测试; 类别 test; 类型 test-coverage; 符号 `_FakeForwardMode`, `is_decode`, `is_extend`) : 测试更新: 为 `fake forward_mode` 添加 `is_decode/is_extend` 方法, 支持新的语法推进逻辑测试。
- python/sglang/srt/speculative/frozen_kv_mtp_worker_v2.py (模块 Frozen MTP; 类别 source; 类型 core-logic; 符号 `forward_batch_generation`) : Frozen KV MTP worker 同步添加 `grammar_barrier` 参数传递, 确保一致性。
- python/sglang/srt/managers/schedule_batch.py (模块 调度批处理; 类别 source; 类型 core-logic) : `ScheduleBatch.copy()` 携带 `has_grammar` 属性, 使队列中的 batch 副本在 `advance_grammar_fsm` 中能够自检是否需要推进语法。

关键符号: `advance_grammar_fsm`, `_advance_pending_grammar`, `run_eagle_verify`, `_allocate_token_bitmask`, `supports_grammar_overlap`, `forward_batch_generation`, `_async_d2h`

关键源码片段

python/sglang/srt/managers/scheduler_components/batch_result_processor.py

核心实现: 新增 `advance_grammar_fsm` 统一 FSM 推进, `_apply_prefill_grammar` 支持 `already_advanced` 标记, 避免重叠路径下重复推进。同时 `_resolve_spec_v2_tokens` 中调用 `advance_grammar_fsm` 并利用 `grammar_retained_tokens` 避免二次 FSM 遍历。

关键改动 1: `_apply_prefill_grammar` 增加 `already_advanced` 参数
当 `grammar barrier` 已推进 FSM 时, 只同步 `grammar.finished` 状态

```
def _apply_prefill_grammar(
    self, *, req: Req, next_token_id: int, already_advanced: bool = False
) -> None:
    # 只有尚未被 barrier 推进时才调用 accept_token
    if not already_advanced:
        try:
            req.grammar.accept_token(next_token_id)
        except ValueError as e:
            logger.error(
                f"Grammar accept_token failed for req {req.rid} "
                f"with token {next_token_id}: {e}"
            )
            req.to_finish = FINISH_ABORT()
    # 始终同步 grammar.finished 状态
    req.grammar.finished = req.finished()
```

关键改动 2: `_resolve_spec_v2_tokens` 中调用 `advance_grammar_fsm`
该行确保即使 `grammar barrier` 未在 `verify` 中执行 (如非重叠路径),

```
# 也能在结果处理时推进 FSM; 且利用 grammar_retained_tokens 避免重复推进
self.advance_grammar_fsm(result, batch)
```

python/sglang/srt/managers/scheduler.py

调度器核心: 修改 `is_disable_overlap_for_batch`, 使支持 grammar overlap 的算法不再因 grammar 强制关闭重叠。新增 `_advance_pending_grammar` 作为 grammar barrier, 在 worker verify 之前推进上一个 batch 的 FSM, 实现 CPU 与 GPU 重叠。

```
# 关键改动 1: is_disable_overlap_for_batch 中新增 supports_grammar_overlap 判断
# 支持 overlap 的算法不再因为 grammar 而强制关闭调度重叠
need_grammar_sync = (
    batch
    and not batch.spec_algorithm.is_none()
    and not batch.spec_algorithm.supports_grammar_overlap() # 新增条件
    and batch.has_grammar
    and batch.forward_mode.is_decode()
    and len(self.result_queue) > 0
)
return disable_overlap_for_batch or need_grammar_sync
```

```
# 关键改动 2: _advance_pending_grammar — grammar barrier
# 遍历结果队列中尚未处理的 decode batch, 推进其语法 FSM,
# 使得后续 verify 步骤的 bitmask 能基于最新的 committed tokens 构建
```

```
def _advance_pending_grammar(self):
    """Grammar barrier: 提前推进队列中待处理 batch 的语法 FSM"""
    for prev_batch, prev_result in self.result_queue:
        self.batch_result_processor.advance_grammar_fsm(prev_result, prev_batch)
```

```
# 调用点: run_batch 中传递 grammar_barrier 给 worker
if batch.spec_algorithm.supports_grammar_overlap():
    fwd_kwargs["grammar_barrier"] = self._advance_pending_grammar
```

python/sglang/srt/speculative/eagle_worker_common.py

共享 verify 函数 `run_eagle_verify` 添加 `grammar_barrier` 参数, 实现异步 D2H 拷贝 + CUDA event, 使 grammar CPU 工作与 target verify forward 重叠。同时将 vocab mask 的 H2D 拷贝改为 `non_blocking`, 进一步隐藏延迟。

```
# eagle_worker_common.py :: run_eagle_verify
# 新增 grammar_barrier 参数, 在目标 verify forward 之后、bitmask 之前调用
```

```
# 1. 异步 D2H 拷贝 (使用 pin-memory + CUDA event)
```

```
grammar_copy_done = None
if batch.has_grammar:
    # _async_d2h 返回 pin-memory 张量, non_blocking=True
    retrieve_next_token_cpu = _async_d2h(verify_input.retrieve_next_token)
    retrieve_next_sibling_cpu = _async_d2h(verify_input.retrieve_next_sibling)
    draft_tokens_cpu = _async_d2h(
        verify_input.draft_token.view(verify_input.retrieve_next_token.shape)
```

```

)
# 记录 event, 表示拷贝发射完成
grammar_copy_done = torch.get_device_module(device).Event()
grammar_copy_done.record()

# ... 发射 target verify forward ...

# 2. 构建 bitmask 前调用 grammar barrier 并同步拷贝
if batch.has_grammar:
    if grammar_barrier is not None:
        grammar_barrier() # 推进上一个 batch 的 FSM
    grammar_copy_done.synchronize() # 等待异步拷贝完成
    vocab_mask = generate_token_bitmask(
        batch.reqs, retrieve_next_token_cpu, retrieve_next_sibling_cpu, draft_tokens_cpu
    )
    if vocab_mask is not None:
        # non_blocking H2D, 真正异步 (因源是 pin-memory)
        vocab_mask = vocab_mask.to(verify_input.retrieve_next_token.device, non_blocking=True)

```

评论区精华

无实质性 review 讨论, PR 由 hnyls2002 直接批准。PR 作者在 body 中详细说明了设计决策和风险, 并请 reviewer 特别关注 EAGLE + grammar 路径的并发正确性。

- 暂无高价值评论线程

风险与影响

- 风险:
 1. 并发正确性: 异步 D2H 拷贝 + CUDA event 的时序依赖较微妙, 如果 event 同步或 grammar barrier 调用时机错误, 可能导致 bitmask 使用未就绪的 draft tokens。PR 作者在 Notes 中明确此路径未在 GPU 充分测试, 需要 reviewer 验证。
 2. Pinned memory 开销: `_allocate_token_bitmask` 改为固定 pin-memory 分配, 可能增加主机内存压力。
 3. 兼容性: 虽然 `supports_grammar_overlap()` 默认返回 False, 其他 spec 算法行为不变, 但若未来新增算法忘记覆盖此方法, 可能意外失去重叠优化 (但不会出错)。
 4. 回归风险: 调度器中 `is_disable_overlap_for_batch` 逻辑调整可能影响非 grammar 场景, 但条件判断仅增加了 `and not batch.spec_algorithm.supports_grammar_overlap()`, 且仅当 `batch.has_grammar` 时才会触发。- 影响: 对用户: grammar constrained + EAGLE decode 的推理速度提升, 吞吐量改善。对系统: 增加少量 pin-memory 使用, CPU 与 GPU 重叠更充分。对团队: 提供了一个可扩展的设计模式, 其他 spec 算法可通过实现 `supports_grammar_overlap()` 获得类似收益。影响范围限定于同时启用 grammar 和 spec decode 的请求, 其他场景无影响。- 风险标记: 并发敏感路径缺少 GPU 验证, pinned memory 额外开销, 旧行为通过默认 False 保留

关联脉络

- 暂无明显关联 PR