

PR #29225 完整报告

sgl-project/sglang

[Spec] Unify spec/non-spec decode result handling and overlap relay-payload gating

合并时间: 2026-06-25 15:44

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/29225>

执行摘要

- 一句话: 统一异构解码结果处理与 relay payload 控制
- 推荐动作: 此 PR 展示了如何在不改变行为的前提下通过抽离公共方法和统一数据结构来降低复杂度, 值得阅读。特别注意 `has_sampled_token_ids` 属性和 `_relay_forward_payload` 的设计。

功能与动机

PR body 指出要 'Unify the overlap relay payload ... and the spec/non-spec decode result handling', 减少重复分支, 提高代码可维护性, 并为后续统一调度铺路。

实现拆解

1. 添加 `has_sampled_token_ids` 属性 (python/sglang/srt/managers/utils.py): 在 `GenerationBatchResult` 中增加 `@property`, 通过 `isinstance(self.next_token_ids, torch.Tensor)` 判断当前迭代是否产生了采样 token, 为后续统一 relay 判断奠定基础。
2. 抽取 `_relay_forward_payload` 方法 (python/sglang/srt/managers/scheduler.py): 将原本分散在 `overlap`、`split-prefill` 和非 `overlap` 路径中的内联 `RelayPayload` 创建与 `self.future_map.stash` 逻辑集中到一个方法中, 并自动处理 `ngram` 跳过。
3. 统一 token 表示 (python/sglang/srt/managers/scheduler_components/batch_result_processor.py): 修改 `_normalize_decode_outputs`, 使 `non-spec` 路径也返回 `List[List[int]]` (即每个请求的单 token 被包在列表中), 后续循环中统一使用 `req.output_ids.extend(next_token_id)` 和 `new_accept_len = len(next_token_id)`, 消除 `if not is_spec` 分支。
4. 统一 `hidden_states` 提取: 将原 `spec` 分支的 `stride = result.speculative_num_draft_tokens` 和 `non-spec` 分支的 `logits_output.hidden_states[i]` 统一为 `stride = result.speculative_num_draft_tokens` 或 `1` 和切片操作 `logits_output.hidden_states[start : start + accept_len]`, 去除了冗余的 `.clone()` 调用。
5. 简化非 `overlap` 中继条件: 在 `scheduler.py` 的 `non-overlap` 路径中, 原本使用 `isinstance(batch_result.next_token_ids, torch.Tensor)` 判断是否需要 `stash`, 现统一使用 `batch_result.has_sampled_token_ids` 属性。

关键文件:

- python/sglang/srt/managers/scheduler.py (模块 调度器; 类别 source; 类型 core-logic; 符号 _relay_forward_payload) : 核心调度器, 新增 _relay_forward_payload 方法统一 relay payload 控制流, 消除多处重复的 RelayPayload 创建和 stash 逻辑。
- python/sglang/srt/managers/scheduler_components/batch_result_processor.py (模块 结果处理; 类别 source; 类型 core-logic; 符号 _normalize_decode_outputs) : 结果处理器, 统一 spec/non-spec 的 token 处理和 hidden_states 提取, 删除大量分支。
- python/sglang/srt/managers/utils.py (模块 数据模型; 类别 source; 类型 core-logic; 符号 has_sampled_token_ids) : 新增 has_sampled_token_ids 属性, 为 relay payload 决策提供统一判断依据。

关键符号: _relay_forward_payload, has_sampled_token_ids

关键源码片段

python/sglang/srt/managers/scheduler.py

核心调度器, 新增 `_relay_forward_payload` 方法统一 relay payload 控制流, 消除多处重复的 RelayPayload 创建和 stash 逻辑。

```
# 在 scheduler.py 中新增方法, 用于统一将当前迭代的 relay payload
# stash 到 FutureMap 中, 供下一次迭代使用。ngram 算法已提前在
# spec_info 中传递 draft, 因此被跳过。
def _relay_forward_payload(
    self,
    future_indices: torch.Tensor,
    batch_result: GenerationBatchResult
) -> None:
    """统一 relay payload 的创建与 stash。"""
    # ngram 不通过 FutureMap 中继, 直接跳过
    if self.spec_algorithm.is_ngram():
        return
    if batch_result.next_draft_input is not None:
        # spec 路径使用 draft_input 构造 RelayPayload
        payload = RelayPayload.from_draft_input(
            batch_result.next_draft_input
        )
    elif batch_result.has_sampled_token_ids:
        # 非 spec 路径使用 next_token_ids 作为 bonus_tokens
        payload = RelayPayload(
            bonus_tokens=batch_result.next_token_ids
        )
    else:
        # 未产生 token (如非最终 PP rank 或非最终 prefill split) ,
        # 无需 stash
        return
    self.future_map.stash(future_indices, payload)
```

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

结果处理器，统一 spec/non-spec 的 token 处理和 hidden_states 提取，删除大量分支。

```
# 在 _normalize_decode_outputs 中，统一输出为 List[List[int]]
def _normalize_decode_outputs(
    self,
    *,
    batch: ScheduleBatch,
    result: GenerationBatchResult,
    logits_output: LogitsProcessorOutput,
    next_token_ids: Union[torch.Tensor, List[int]],
) -> Tuple[Union[List[int], List[List[int]]], Optional[List[float]]]:
    next_token_logprobs = None
    if not batch.spec_algorithm.is_none():
        # spec 路径：解包 validate 输出
        next_token_ids = self._resolve_spec_v2_tokens(result, batch)
    else:
        # non-spec 路径：确保为列表并包装成每请求单 token 的列表
        ids = (
            next_token_ids.tolist()
            if torch.is_tensor(next_token_ids)
            else next_token_ids
        )
        next_token_ids = [[t] for t in ids]
    # ... logprob 处理保持不变
    return next_token_ids, next_token_logprobs

# 在 process_batch_result_decode 循环中，统一处理
for i, req in enumerate(batch.reqs):
    # next_token_id 始终为 list (non-spec 为 [token])
    next_token_id = next_token_ids[i]
    is_spec = not batch.spec_algorithm.is_none()
    # 统一使用 extend, non-spec 追加单元素
    req.output_ids.extend(next_token_id)
    new_accept_len = len(next_token_id)
    # ...
    if req.return_hidden_states and logits_output.hidden_states is not None:
        # 统一 hidden_states 切片，stride 在 non-spec 时为 1
        stride = result.speculative_num_draft_tokens or 1
        accept_len = len(next_token_id)
        start = i * stride
        req.hidden_states.extend(
            logits_output.hidden_states[start : start + accept_len]
            .cpu()
            .tolist() # 已移除原 spec 分支的 .clone()
        )
```

评论区精华

无公开 review 讨论；变更经作者独立审阅后合并。

- 无公开讨论 (other): 无争议, 直接合并。

风险与影响

- 风险: 作为重构, 行为保留, 但需注意: (1) `batch_result_processor.py` 中 `hidden_states` 提取统一为 `stride = result.speculative_num_draft_tokens or 1`, 需确认 non-spec 时 `stride=1` 切片行为与原 `logits_output.hidden_states[i]` 完全一致 (原为单个行索引, 新为切片 `start:start+accept_len`, 当 `accept_len=1` 时效果相同但涉及 `tolist()` 方式可能不同; 原用法有 `.clone()` 新去掉了, 需确保共享内存不影响)。 (2) relay payload 判断逻辑从调用点内联改为集中方法, 需确认 ngram 和分 Prefill Split 场景正确性。 (3) 缺少直接对应的测试文件变更, 回归风险存在。
- 影响: 对用户无功能影响; 对系统减少了代码重复和分支判断; 对团队提升了调度器与结果处理器的可维护性。
- 风险标记: 缺少测试覆盖, 核心路径重构, `hidden_states` 切片行为变化需验证

关联脉络

- PR #29220 [Spec] Dissolve EagleDraftInputV2Mixin so spec-info dataclasses hold data only: 同为 speculative-decoding 调度 / 数据模型重构, 延续统一 spec/non-spec 处理的趋势。
- PR #27625 Remove `Req.extend_logprob_start_len` field and make it pure: 也是移除调度器中分支、统一数据结构的重构。