

PR #28500 完整报告

sgl-project/sglang

[Perf] Make spec-decode penalty H2D non-blocking and share decode cumulate path

合并时间: 2026-06-17 15:59

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/28500>

执行摘要

- 一句话: 规格化解码 penalty H2D 异步 + 去重
- 推荐动作: 值得精读, 特别是学习如何通过 `pin_memory + non_blocking=True` 实现异步 H2D 拷贝的模式。设计决策清晰: 共享方法减少重复, 性能提升明确。

功能与动机

PR #28491 已修复非 spec 路径的 blocking H2D 问题, 但 spec-decode 路径 (`EagleDraftInput.prepare_for_decode`) 仍然使用阻塞的 `torch.tensor(..., device=...)`, 在 overlap scheduler 下导致调度流每步同步。本 PR 旨在消除 spec-path stall 并消除代码重复。

实现拆解

1. 提取共享方法 `cumulate_penalty_output_tokens` 到 `ScheduleBatch`: 将原来内联在 `prepare_for_decode` 中的 penalty 累积逻辑提取为独立方法, 添加 `pin_memory` 确保异步 H2D 真正生效。
2. 非 spec 路径(`schedule_batch.py`): `prepare_for_decode` 中直接调用 `self.cumulate_penalty_output_tokens()`, 替换原来内联代码。
3. spec 路径(`eagle_info_v2.py`): `EagleDraftInputV2Mixin.prepare_for_decode` 中同样调用 `batch.cumulate_penalty_output_tokens()`, 删除原有的 `torch.tensor(..., device=batch.device)` 内联实现。
4. 提交演进: 第一 commit 实现去重和非阻塞; 第二 commit 根据 review 建议添加 `pin_memory` 参数确保异步拷贝真正生效。

关键文件:

- `python/sglang/srt/managers/schedule_batch.py` (模块 调度器; 类别 source; 类型 core-logic; 符号 `cumulate_penalty_output_tokens`): 提取了 `cumulate_penalty_output_tokens` 方法, 并修改 `prepare_for_decode` 调用它; 核心逻辑和异步优化在此实现。
- `python/sglang/srt/speculative/eagle_info_v2.py` (模块 推测解码; 类别 source; 类型 core-logic): spec-decode 路径从内联改为调用 `batch.cumulate_penalty_output_tokens()`, 消除重复并享受异步优化。

关键符号: `cumulate_penalty_output_tokens`, `prepare_for_decode`

关键源码片段

python/sglang/srt/managers/schedule_batch.py

提取了 `cumulate_penalty_output_tokens` 方法，并修改 `prepare_for_decode` 调用它；核心逻辑和异步优化在此实现。

```
def cumulate_penalty_output_tokens(self):
    # Under overlap batch.input_ids is just a placeholder here -- the
    # real token is relayed via future_map and resolved at forward
    # entry. So take the last output token from Req directly
    # (origin_input_ids[-1] on the first decode, before any output).
    last_tokens = [
        req.output_ids[-1] if len(req.output_ids) else req.origin_input_ids[-1]
        for req in self.reqs
    ]
    # Non-blocking H2D so this per-step copy doesn't sync behind the forward.
    # pin_memory (matching the prefill-path tensors) keeps the copy async;
    # is_pin_memory_available falls back to pageable on unsupported devices.
    latest_output_ids = torch.tensor(
        last_tokens,
        dtype=torch.int64,
        pin_memory=is_pin_memory_available(self.device),
    ).to(self.device, non_blocking=True)
    self.sampling_info.penalizer_orchestrator.cumulate_output_tokens(
        latest_output_ids
    )
```

评论区精华

review 中 `gemini-code-assist[bot]` 指出：要使 `non_blocking=True` 真正异步，CPU tensor 必须 reside in pinned memory，否则 PyTorch 会回退到同步拷贝。建议使用 `pin_memory=is_pin_memory_available(self.device)`。该建议已被采纳在第二 commit 中。

- 使用 `pin_memory` 确保非阻塞 H2D 拷贝真正异步 (performance): 采纳建议，在 `torch.tensor` 中添加 `pin_memory=is_pin_memory_available(self.device)`。

风险与影响

- 风险：风险较低：变更仅为提取共享方法和添加 `pin_memory` 参数，逻辑不变。
`pin_memory` 在设备不支持时通过 `is_pin_memory_available` 安全回退。未发现测试文件变更，但 CI 运行了 `test_penalty.py` 和 `test_spec_eagle.py` 并通过。
- 影响：影响范围：规格化解码场景下 `penalty` 开启时的 `decode` 性能。消除 per-step 同步后，`overlap scheduler` 下的 `decode` 吞吐应有提升。对非 `spec` 路径无功能影响。
- 风险标记：异步 H2D 正确性依赖 pinned memory

关联脉络

- PR #28491 [Perf] Make latest_output_ids H2D non-blocking in prepare_for_decode: 本 PR 是 #28491 的 follow-up, 将其非 spec 路径的异步 H2D 优化扩展到 spec-decode 路径, 并提取共享方法。