

PR #51650 完整报告

vllm-project/vllm

[PP][XPU]Overlap async-scheduling PP sampled-token broadcast with compute

合并时间: 2026-08-14 17:00

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/51650>

执行摘要

- 一句话: PP 下 sampled-token 广播改异步, XPU 吞吐提升约 33%
- 推荐动作: 值得精读, 虽然改动只有 11 行, 但展示了「把阻塞通信改为异步 + 延迟 wait 到消费点」的经典重叠手法, 并暴露了 Work 句柄生命周期管理这一易错点。建议关注 `_pp_recv_work` 的消费时序, 并在未来补一个 PP + async 的时序测试。

功能与动机

PR body 明确指出: 当 async scheduling 与 PP 同时启用时, 末 stage 产生的 sampled token ids 通过 `torch.distributed.broadcast` 回传首 stage, 原实现是 "a blocking collective (`async_op=False`) on the default compute stream", 阻塞了计算流。目的就是让 "All token-independent work now can overlaps", 即把通信与计算重叠起来, 减少每轮迭代中集体通信对流水线的停顿。

实现拆解

1. 新增状态字段: 在 `vllm/v1/worker/gpu_model_runner.py` 的 `__init__` 中新增 `self._pp_recv_work: torch.distributed.Work | None = None`, 用于保存尚未完成的异步广播句柄, 初始化时为空。
2. 接收侧改为异步广播: 在 `_pp_receive_prev_sampled_token_ids_to_input_batch` 中, 将原来的阻塞式 `torch.distributed.broadcast(recv, src=pp.last_rank, group=pp.device_group)` 改为 `async_op=True`, 并把返回的 `Work` 对象赋给 `self._pp_recv_work`; chunked prefill 分支依旧跳过广播, 此时句柄保持为 `None`。
3. 消费前统一 wait: 在 `_prepare_input_ids` 入口处、读取 `prev_sampled_token_ids` 之前, 执行 `if self._pp_recv_work is not None: self._pp_recv_work.wait(); self._pp_recv_work = None`, 将阻塞点推迟到真正消费 tokens 的前一刻, 使调度准备、输入拷贝等 token 无关工作与通信重叠。
4. 测试与配置配套: 本次没有新增测试文件或配置项; 性能验证在 XPU 设备上用 `gpt-oss-20B + sonnet` 数据集完成 (`pp=2, tp=2`), 吞吐从 338 tok/s 提升到 448 tok/s。

关键文件:

- `vllm/v1/worker/gpu_model_runner.py` (模块 模型执行; 类别 source; 类型 core-logic; 符号 `init, _prepare_input_ids, _pp_receive_prev_sampled_token_ids_to_input_batch`): 唯一变更文件, 完成 PP 下 sampled-token 广播从阻塞到异步的改造, 并新增

_pp_recv_work 句柄管理。

关键符号：init, _prepare_input_ids, _pp_receive_prev_sampled_token_ids_to_input_batch

关键源码片段

vllm/v1/worker/gpu_model_runner.py

唯一变更文件，完成 PP 下 sampled-token 广播从阻塞到异步的改造，并新增 _pp_recv_work 句柄管理。

```
# vllm/v1/worker/gpu_model_runner.py (精简整理)
```

```
# 核心思路：把 PP 下 sampled-token 广播从阻塞改为异步，把 wait() 推迟到真正消费 token 之前，  
# 让调度、输入准备等 token 无关工作可以与通信重叠。
```

```
def _prepare_input_ids(self, scheduler_output, num_reqs, total_num_scheduled_tokens, cu_num_tokens) -> None:
```

```
    # 先同步上一轮发起的异步 PP broadcast，确保 sampled tokens 已经就绪。
```

```
    # 只有在存在未完成的 Work 句柄时才 wait；chunked prefill 跳过广播时 _pp_recv_work 为 None。
```

```
    if self._pp_recv_work is not None:
```

```
        self._pp_recv_work.wait()
```

```
        self._pp_recv_work = None
```

```
    if self.input_batch.prev_sampled_token_ids is None:
```

```
        # 正常调度路径：直接拷贝 CPU 侧 input_ids。
```

```
        self.input_ids.copy_to_gpu(total_num_scheduled_tokens)
```

```
        if self.enable_prompt_embeds:
```

```
            self.inputs_embeds.copy_to_gpu(total_num_scheduled_tokens)
```

```
            self.is_token_ids.copy_to_gpu(total_num_scheduled_tokens)
```

```
        return
```

```
    # ... 后续按 prev_positions 将 prev_sampled_token_ids 拷贝到 input_ids 对应槽位（略）
```

```
# —— 接收侧：非最后 PP stage 接收 sampled tokens 并异步发起广播 ——
```

```
def _pp_receive_prev_sampled_token_ids_to_input_batch(self) -> None:
```

```
    pp = get_pp_group()
```

```
    assert not pp.is_last_rank
```

```
    num_reqs = self.input_batch.num_reqs
```

```
    # prev_sampled_token_ids 形状固定为 [num_reqs, 1]。
```

```
    recv = torch.empty((num_reqs, 1), dtype=torch.int32, device=self.device)
```

```
    # chunked prefill 时 sampled tokens 是占位值且会被丢弃，无需广播。
```

```
    if not self._is_all_reqs_chunked_prefill():
```

```
        # 关键变更：async_op=True 让 broadcast 在后台执行，返回 Work 句柄，
```

```
        # 由下一轮 _prepare_input_ids 开头统一 wait()，从而与计算重叠。
```

```
        self._pp_recv_work = torch.distributed.broadcast(
```

```
            recv, src=pp.last_rank, group=pp.device_group, async_op=True
```

```
        )
```

```
    self.input_batch.prev_sampled_token_ids = recv
```

```
    # 构建 prev_req_id_to_index 映射，供 _prepare_input_ids 做 req_id -> 上一批行的映射；
```

```
    # 同时按需求追加 -1 占位 token 并推进 num_tokens_no_spec（此部分为逐请求簿记，可省略）。
```

评论区精华

review 中无实质技术争议: jikunshang 批准并留言 "LGTM. @njhill can you help take a review? thanks!", rogerxfeng8 也给出 approved; claude[bot] 提示这是 fork PR, 自动 review 被禁用。整体属于小范围、低争议的异步化优化, 直接合入。

- Review 与后续复审请求 (other): 无实质技术争议, 改动合入 main。

风险与影响

- 风险:

1. 异步句柄生命周期: `_pp_recv_work` 依赖 `_prepare_input_ids` 在下一轮广播发起前被调用并消费该句柄; 若未来某条路径跳过 `_prepare_input_ids` 或重复调用接收函数, 旧句柄可能被覆盖, 产生隐式等待或悬挂。
2. 后端差异: `async_op=True` 的语义在 NCCL、OneCCL 等不同后端上可能不同, 本 PR 仅在 XPU 实测; NVIDIA 等平台未验证, 存在流同步行为差异的回归风险。
3. 测试覆盖缺失: 没有针对 PP + async scheduling 时序的单元测试, 回归只能依赖现有 CI, 风险集中在 `_pp_receive_prev_sampled_token_ids_to_input_batch` 与 `_prepare_input_ids` 的契约上。
4. chunked prefill 分支: `_is_all_reqs_chunked_prefill()` 为真时跳过广播, `_pp_recv_work` 保持 None, 逻辑正确, 但与新增 wait 路径的交互未被测试显式覆盖。
 - 影响: 影响范围限定在 `async_scheduling=True` 且 `pp>1` 的非末 stage 推理路径: 吞吐收益在 XPU 上显著 (约 +32.5%), 对其他后端理论上同样受益但未验证; 无 API、配置、模型兼容性变化, 用户无感。对团队而言, 这是 PP + async scheduling 通信路径上的样板优化, 后续可推广到 draft token、logits 等其他集体通信。
 - 风险标记: PP+async 核心路径变更, 缺少直接测试覆盖, 仅 XPU 实测验证, 异步句柄生命周期依赖调用顺序

关联脉络

- PR #52329 [Performance][MRV2] Cache logits-processing request state: 同为 v1 worker 采样路径的性能优化, 与 PP+async 下采样结果处理的时序优化属同一 MRV2 性能演进脉络。
- PR #52482 [Bugfix][V1][Multimodal] Ignore stale same-step encoder cache evictions: 同为 async scheduling 下迭代间时序问题的处理, 侧面说明该调度模式对迭代间状态一致性的敏感性。