

PR #7041 完整报告

verl-project/verl

[fully_async] fix: qwen2.5-0.5b fully async OOM bug fix

合并时间: 2026-07-15 09:28

原文链接: <http://prhub.com.cn/verl-project/verl/pull/7041>

执行摘要

- 一句话: 降低权重同步桶大小和 GPU 显存占比, 修复 OOM
- 推荐动作: 该 PR 是典型的运行时显存调优, 无复杂逻辑变更, 可作为全异步训练 OOM 问题的参考案例; 若你也在 Ascend 上跑全异步训练, 建议同步此调整。

功能与动机

PR body 指出 `recv_buf` 是全异步算法中权重同步需要的额外内存, 之前分配过大或推理引擎占用过多 GPU 显存导致 OOM, 需减小 `recv_buf` 大小和显存占用比例。

实现拆解

1. 在测试脚本 `tests/special_e2e/run_fully_async_policy.sh` 的 `fsdp2` 分支中, 将 `actor_rollout_ref.rollout.checkpoint_engine.update_weights_bucket_megabytes` 从 1024 改为 512, 减小单个权重同步桶的大小, 使 `recv_buf` 分配更轻量。
2. 在同一脚本的 NPU 特定配置中, 将 `actor_rollout_ref.rollout.gpu_memory_utilization` 从 0.70 降为 0.50, 降低推理引擎的显存占用比例, 为 `recv_buf` 和其他训练组件腾出空间。
3. 两项变更均针对 `fsdp2` 策略下的全异步训练场景, 属于显存使用量的调优修复。

关键文件:

- `tests/special_e2e/run_fully_async_policy.sh` (模块 全异步策略; 类别 test; 类型 test-coverage): 唯一变更文件, 通过降低权重同步桶大小和 GPU 显存利用率来修复 OOM。

关键符号: 未识别

评论区精华

无 review 讨论, 仅 robot 评论总结变更内容。

- 暂无高价值评论线程

风险与影响

- 风险: 风险较低: 仅涉及测试脚本中的两个配置参数调整, 极小幅度可能影响训练速度 (更小的桶增加通信次数, 更低的显存利用率可能降低推理吞吐), 但主要目标是修复 OOM,

不会引入功能性破坏。

- 影响：直接影响 fsdp2 全异步训练在 Ascend NPU 上的显存使用，修复了 OOM 问题；对其他训练模式（megatron、非全异步）无影响；用户若需要复原旧配置可自行调整脚本。
- 风险标记：配置参数调整，主要影响 NPU 场景

关联脉络

- PR #6897 [tool, rollout] feat: Adapt SkipManager on Trainer V1: 同为全异步训练相关，涉及 SkipManager 适配 Trainer V1，本 PR 修复了该功能的显存瓶颈。
- PR #7032 [tool, rollout] feat: support parameter sync steps in Skip Manager: 全异步训练参数同步优化，本 PR 的权重同步桶调整属于参数同步的配套调优。