

PR #7049 完整报告

verl-project/verl

[trainer, perf] fix: include standalone rollout GPUs in throughput denominator for separate async

合并时间: 2026-07-20 11:33

原文链接: <http://prhub.com.cn/verl-project/verl/pull/7049>

执行摘要

本 PR 修复了分离异步 (separate-async) 模式下吞吐量指标计算不准确的问题。之前仅使用 trainer 端 GPU 数作为分母, 忽略了独立的 rollout GPU, 导致吞吐量被高估。通过在基类新增 `_get_n_gpus_for_throughput` 方法并让分离异步 trainer 重写以包含 rollout GPU 数, 确保吞吐量计算准确且与同步 / 同址异步模式可比。

功能与动机

分离异步模式下, rollout 运行在 trainer 资源池之外的专用 GPU 上。原代码中 `compute_throughout_metrics` 使用 `resource_pool_manager.get_n_gpus()` 作为分母, 只统计了 trainer 端 GPU, 导致吞吐量 (throughput) 被夸大, 无法与同步模式和同址异步模式进行公平比较。

实现拆解

1. 基类 `trainer_base.py` 新增 `_get_n_gpus_for_throughput` 方法: 默认返回 trainer 端 GPU 数, 与原有逻辑一致。同时在 `_compute_metrics` 中将直接调用 `self.resource_pool_manager.get_n_gpus()` 替换为 `self._get_n_gpus_for_throughput()`, 使子类可通过重写影响结果。
2. 分离异步 trainer `trainer_separate_async.py` 重写上述方法: 返回 `trainer_gpus + rollout_gpus`, 其中 rollout GPU 数通过配置 `actor_rollout_ref.rollout.n_gpus_per_node * actor_rollout_ref.rollout.nnodes` 计算。
3. 测试配套: 未新增测试, 依赖现有 CI 覆盖。

见 `key_files` 中的 `annotated_snippet_markdown`。

评论区精华

无人工审核评论。审核者 `wuxibin89` 已批准。

风险与影响

- 风险: 低。变更范围小, 默认行为无变化。若将来有新的子类未重写该方法, 其吞吐量分母仍为 trainer GPU 数 (与之前一致), 无回归风险。
- 影响: 仅影响分离异步模式的吞吐量指标, 使其更准确。用户无需配置修改。

关联脉络

与近期 PR #7041 (fully async OOM 修复) 同属异步模式优化系列，表明团队正持续打磨异步训练的性能指标和稳定性。