

PR #7595 完整报告

verl-project/verl

[trainer] feat: Update run_qwen3_30b_veomni configuration

合并时间: 2026-08-29 09:39

原文链接: <http://prhub.com.cn/verl-project/verl/pull/7595>

执行摘要

- 一句话: 更新 30B veomni 脚本的专家并行配置
- 推荐动作: 该 PR 为低优先级的配置调整, 不值得精读。可留意其将专家并行大小参数化并默认调整为 8 的做法, 这反映了 veomni 在 NPU 上对专家并行的偏好。

功能与动机

PR body 中说明: 'If the EP1 is used, the performance will deteriorate.' 即使用专家并行大小 1 (EP=1) 会导致性能下降, 因此需要将专家并行大小配置为可调且默认使用更大的值 (8) 以提升性能。

实现拆解

1. 在 examples/grpo_trainer/run_qwen3_30b_veomni.sh 中, 于 NPU 分支添加了 `ep_size=${ep_size:-8}` 变量定义, 允许通过环境变量覆盖, 默认值为 8。
2. 将 `actor_rollout_ref.actor.veomni.expert_parallel_size=1` 修改为 `actor_rollout_ref.actor.veomni.expert_parallel_size=${ep_size}`, 使得专家并行大小可配置。
3. 将 `VLLM_VERSION=0.13.0` 更新为 `VLLM_VERSION=0.23.0`, 以使用较新版本, 可能期望获得更好的性能或特性。
4. 变更仅涉及脚本配置, 无配套测试或文档更新。

关键文件:

- examples/grpo_trainer/run_qwen3_30b_veomni.sh (模块 示例脚本; 类别 other; 类型 core-logic): 唯一变更文件, 调整专家并行大小和 VLLM 版本配置

关键符号: 未识别

关键源码片段

`examples/grpo_trainer/run_qwen3_30b_veomni.sh`

唯一变更文件, 调整专家并行大小和 VLLM 版本配置

```
# 以下是 run_qwen3_30b_veomni.sh 中的关键片段 (已整理)
# 在 NPU 设备分支中, 新增 ep_size 变量, 默认值为 8
ep_size=${ep_size:-8} # 通过环境变量可覆盖默认 EP 大小
```

```
# 在 ACTOR 配置中，将专家并行大小改为动态获取 ep_size
actor_rollout_ref.actor.veomni.expert_parallel_size=${ep_size} # 使用 EP=
8（默认）可避免性能下降
```

```
# 同时升级 vLLM 版本至 0.23.0
VLLM_VERSION=0.23.0
```

评论区精华

该 PR 没有评论或审阅评论。审阅者 wucong25 直接批准，无讨论。

- 暂无高价值评论线程

风险与影响

- 风险：风险较低。变更仅涉及示例脚本，不会影响核心代码路径。主要风险是：
 - 默认 EP=8 可能在某些环境下（如设备数不足或不支持）导致启动失败或配置不匹配，但脚本中提供了覆盖机制，可通过设置 ep_size 调整。
 - VLLM 版本升级从 0.13.0 到 0.23.0 可能引入不兼容的 API 变更，但仅限于脚本环境配置，不影响整体框架。
 - 未提供测试验证，属于示例脚本，风险可控。
- 影响：影响范围有限：仅影响使用该脚本的用户，主要是 NPU 环境下运行 30B veomni GRPO 训练的用户。预期提升训练性能（通过使用更大的专家并行度）。对系统其他部分无影响。
- 风险标记：示例脚本变更，无测试覆盖

关联脉络

- PR #7577 [veomni] feat: DeepSeek V4 QAT bf16 fake quant training: 同样涉及 veomni 训练脚本配置，表明 veomni 相关示例脚本的演进方向