

PR #7629 完整报告

verl-project/verl

[ci] chore: fix ci failure

合并时间: 2026-08-31 14:18

原文链接: <http://prhub.com.cn/verl-project/verl/pull/7629>

执行摘要

- 一句话: 修复 CI 失败: E2E 补 max_model_len 并注释不稳定测试
- 推荐动作: 值得快速浏览, 不必精读。建议与 PR#7632 成对阅读, 理解 verl 团队应对 vLLM chunked prefill 校验收紧的两层策略; 同时留意 vllm.yml 中被注释的 determinism 测试是否有恢复计划。

功能与动机

PR body 仅说明 "As title.", 具体动机从代码注释还原: vLLM 拒绝

enable_chunked_prefill=False 时 max_num_batched_tokens < max_model_len 的组合; Qwen2.5 / Qwen3-VL 默认 max_model_len 取 max_position_embeddings (32k / 128k), 远超 E2E 实际 prompt + response 预算, 导致 CI 启动崩溃。此外 vllm.yml 的生成确定性测试持续不稳定, 也被注释掉。

实现拆解

1. tests/special_e2e/run_ppo_trainer_veomni.sh: 把硬编码的 data.max_prompt_length=512 / data.max_response_length=128 改为 MAX_PROMPT_LEN / MAX_RESPONSE_LEN 环境变量, 新增 MAX_MODEL_LEN (默认二者之和 640), 并在 common_params 中追加 actor_rollout_ref.rollout.max_model_len。这样 vLLM 的 max_num_batched_tokens 与 max_model_len 的关系满足校验, 启动不再崩溃。
2. tests/special_e2e/ppo_trainer/run_function_reward.sh: 在变量区追加 MAX_MODEL_LEN (默认 1024), 并在 rollout 参数中传入 max_model_len。注意该脚本默认 enable_chunked_prefill=True, 但同样为 VL 模型关闭 prefill 的场景做好准备。
3. .github/workflows/vllm.yml: 将 "Test vllm rollout generation determinism" 整个 job 注释掉, 让步 CI 通过; 这是本次修复中唯一以牺牲测试覆盖为代价的改动。
4. .github/workflows/e2e_ppo_trainer_veomni_vllm.yml: 删除 install 阶段重复的 pip3 install -r requirements.txt (保留 requirements-test.txt), 移除网络诊断注释, 并把 cleanup job 的 needs 列表压缩为单行。配套: 无新增单元测试; E2E 脚本正常运行即为验证手段。新增 MAX_MODEL_LEN 等环境变量入口, 作为脚本配置配套。

关键文件:

- tests/special_e2e/run_ppo_trainer_veomni.sh (模块 端到端测试; 类别 test; 类型 test-coverage) : 核心修复: 推导 MAX_MODEL_LEN 并传给 vLLM, 规避 enable_chunked_prefill=False 时 max_num_batched_tokens 校验失败导致的启动崩溃。
- tests/special_e2e/ppo_trainer/run_function_reward.sh (模块 端到端测试; 类别 test; 类型 test-coverage) : 同一问题在 function_reward E2E 的同步修复, 确保 VL 模型 (默认 128k) 下关闭 chunked prefill 不崩溃。
- .github/workflows/vllm.yml (模块 持续集成; 类别 infra; 类型 infrastructure) : 本次 CI 修复的直接手段之一: 注释掉经常失败的 vLLM 生成确定性测试 job, 让 CI 恢复绿色; 代价是确定性覆盖暂时缺失。
- .github/workflows/e2e_ppo_trainer_veomni_vllm.yml (模块 持续集成; 类别 infra; 类型 infrastructure) : 清理 install 阶段多余的 requirements.txt 安装和诊断注释, 并格式化 cleanup job 依赖, 减少 CI 不稳定因素。

关键符号: 未识别

关键源码片段

tests/special_e2e/run_ppo_trainer_veomni.sh

核心修复: 推导 MAX_MODEL_LEN 并传给 vLLM, 规避 enable_chunked_prefill=False 时 max_num_batched_tokens 校验失败导致的启动崩溃。

```
# vLLM 在关闭 chunked prefill 时会校验 max_num_batched_tokens >= max_model_len,
# 不满足则拒绝启动。Qwen2.5 / Qwen3-VL 默认 max_model_len 取
# max_position_embeddings (32k / 128k), 远大于 E2E 实际预算, 因此显式 pin 到
# prompt + response 之和, 既满足校验又避免预留过多 KV cache。
MAX_PROMPT_LEN=${MAX_PROMPT_LEN:-512}
MAX_RESPONSE_LEN=${MAX_RESPONSE_LEN:-128}
MAX_MODEL_LEN=${MAX_MODEL_LEN:-$((MAX_PROMPT_LEN + MAX_RESPONSE_LEN))}

# common_params 数组内对应调整: 硬编码的 512 / 128 换成变量引用,
# 并在 enable_chunked_prefill=False 之后追加 max_model_len 参数。
common_params=(
    ...
    data.max_prompt_length="${MAX_PROMPT_LEN}" \
    data.max_response_length="${MAX_RESPONSE_LEN}" \
    ...
    actor_rollout_ref.rollout.enable_chunked_prefill=False \
    actor_rollout_ref.rollout.max_model_len="${MAX_MODEL_LEN}" \
    ...
)
```

tests/special_e2e/ppo_trainer/run_function_reward.sh

同一问题在 function_reward E2E 的同步修复, 确保 VL 模型 (默认 128k) 下关闭 chunked prefill 不崩溃。

```
# VL 模型默认 max_model_len 可达 128k, 本脚本若关闭 chunked prefill,
# vLLM 会因 max_num_batched_tokens < max_model_len 拒绝启动。
```

```
# 这里把 max_model_len 固定到实际 prompt + response 预算（默认 1024）。
MAX_MODEL_LEN=${MAX_MODEL_LEN:-$((MAX_PROMPT_LEN + MAX_RESPONSE_LEN))}

# common_params 中的 rollout 参数片段:
actor_rollout_ref.rollout.enable_chunked_prefill="${ENABLE_CHUNKED_PREFILL}" \
actor_rollout_ref.rollout.max_model_len="${MAX_MODEL_LEN}" \
```

评论区精华

本 PR 没有任何 review 评论或 issue 讨论 (comments_count = 0, review_comments_count = 0)，由作者 wuxibin89 自行合入。因此没有可提炼的公开交锋；设计权衡只能从代码注释与 diff 中推断：最大化 CI 稳定性优先于保留确定性测试覆盖。

- 暂无高价值评论线程

风险与影响

- 风险：风险：1) 测试覆盖缺口：vllm.yml 中 determinism 测试被注释，vLLM 生成确定性回归暂时不设防，需在根因修复后恢复。2) 配置一致性：MAX_MODEL_LEN 默认由 MAX_PROMPT_LEN + MAX_RESPONSE_LEN 推导，如果外部只覆写其中一个变量，可能导致推导值偏离实际数据预算，需要同步设置。3) 依赖完整性：veomni_vllm workflow 删除 requirements.txt 安装，若 requirements-test.txt 未完整包含运行依赖会引入缺失（通常已覆盖，但值得确认）。4) 修复范围有限：其他 E2E 脚本若同样以 enable_chunked_prefill=False 运行，仍可能遇到相同启动失败。
- 影响：影响范围仅涉及 CI 工作流与 E2E 测试脚本：vllm.yml 的 vLLM 生成确定性测试暂时下线，veomni E2E 恢复通过；对用户训练 / 推理路径无运行时影响。对团队而言，CI 稳定性优先于覆盖完整性的取舍值得记录，后续需要跟进恢复 determinism 测试并落实根本修复。
- 风险标记：测试覆盖缺失，配置一致性风险，修复范围有限

关联脉络

- PR #7632 [vllm] fix: raise max_num_batched_tokens to max_model_len when chunked prefill is disabled: 同一 vLLM 启动校验问题的根本修复：在 vllm_async_server.py 中提高 max_num_batched_tokens；本 PR 是 E2E 脚本侧的配套规避。
- PR #7584 [ci] fix: drop stale enable_chunked_prefill=False from Ascend NPU scripts: 同样围绕 enable_chunked_prefill=False 与 vLLM 行为变化的 CI 清理，属于同一条修复线索。