

PR #42042 完整报告

vllm-project/vllm

[CI][Examples][RLHF] Disable async scheduling in rlhf_async_new_apis

合并时间: 2026-05-08 19:58

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/42042>

执行摘要

- 一句话: RLHF 示例禁用异步调度以规避 Bug
- 推荐动作: 该 PR 适合快速合入以修复 CI 稳定性问题, 技术洞察价值有限。建议关注关联 Issue #42043 的根本修复, 届时移除该工作区。

功能与动机

RayExecutorV2 成为默认后端后, CI 中 `examples/rl/rlhf_async_new_apis.py` 在 2 GPU H100 测试中持续失败。根本原因是 `AsyncScheduler` 在 `pause_generation(mode="keep")` 缓存重置后丢弃了第一个 post-resume token, 使得训练后权重同步验证阶段的确定性比对无法通过。PR body 明确指出该工作区需要等待 #42043 的修复。

实现拆解

1. 定位入口: 修改 `examples/rl/rlhf_async_new_apis.py` 中两个 `llm_kwargs` 字典的构建位置。
2. 首次添加: 在 `llm_kwargs` (第 222 行) 中增加 `async_scheduling=False`, 并附上 TODO 注释, 引用修复 Issue #42043, 强调两个实例必须一致。
3. 二次添加: 根据 Code Review 反馈, 在 `llm_v2_kwargs` (第 363 行) 中增加相同的 `async_scheduling=False` 和 TODO 注释, 确保两个 LLM 实例配置对称, 避免后续维护陷阱。
4. 无其他变更: 仅添加 6 行代码 (含注释), 无测试、配置或部署文件改动。

关键文件:

- `examples/rl/rlhf_async_new_apis.py` (模块 RLHF 示例; 类别 source; 类型 entrypoint) : 唯一的变更文件, 是 RLHF 异步 API 的 CI 示例入口脚本。PR 在此文件中为两个 LLM 实例添加 `async_scheduling=False` 以规避 `AsyncScheduler` 的 token 丢失 Bug。

关键符号: 未识别

关键源码片段

`examples/rl/rlhf_async_new_apis.py`

唯一的变更文件, 是 RLHF 异步 API 的 CI 示例入口脚本。PR 在此文件中为两个 LLM 实例添加 `async_scheduling=False` 以规避 `AsyncScheduler` 的 token 丢失 Bug。

```

# 第一个 LLM 实例 (V1) 的配置参数字典
llm_kwargs = dict(
    model=MODEL_NAME_V1,
    enforce_eager=True,
    max_model_len=8192,
    distributed_executor_backend="ray",
    attention_backend=ATTN_BACKEND,
    gpu_memory_utilization=0.75,
    weight_transfer_config=WeightTransferConfig(backend="nccl"),
    # 临时禁用异步调度, 直到 #42043 被修复。
    # 两个 LLM 实例必须保持一致的配置。
    async_scheduling=False,
)
llm_kwargs.update(rocm_determinism_kwargs)

# ... (中间代码省略) ...

# 第二个 LLM 实例 (V2) 的配置参数字典
llm_v2_kwargs = dict(
    model=MODEL_NAME_V2,
    enforce_eager=True,
    max_model_len=8192,
    gpu_memory_utilization=0.75,
    distributed_executor_backend="ray",
    attention_backend=ATTN_BACKEND,
    # 临时禁用异步调度, 直到 #42043 被修复。
    # 两个 LLM 实例必须保持一致的配置。
    async_scheduling=False,
)
llm_v2_kwargs.update(rocm_determinism_kwargs)

```

评论区精华

Gemini Code Assist Bot 在 Review 中提出两个实例需要对称的 **TODO** 注释, 确保后续修复时统一跟踪。作者采纳并补充了第二个实例的注释。NickLucche批准了该PR。无其他争议讨论。

- 为第二个 LLM 实例补充 TODO 注释 (design): 作者接受了建议, 在第二个实例中添加了相同的注释。

风险与影响

- 风险: 低风险。变更仅涉及示例脚本中的配置参数, 不改变核心引擎逻辑。禁用 `async_scheduling` 可能会轻微影响异步场景下的吞吐量, 但 RLHF 示例主要用于验证权重同步的正确性, 而非性能基准。后续需同步跟进 #42043 修复以重新启用异步调度。
- 影响: 影响范围局限于 `examples/rl/rlhf_async_new_apis.py` 这一个示例文件。对用户而言, RLHF 示例现在能稳定通过 CI 验证阶段; 对系统无影响; 对团队降低了 CI 误报率。影响程度低。
- 风险标记: 临时工作区, 需后续跟进

关联脉络

- PR #41421 [Core] Make RayExecutorV2 the default executor backend: 引入 RayExecutorV2 为默认后端，直接导致 AsyncScheduler 行为变化，触发本 PR 需要规避的 Bug。
- PR #42043：该 Issue 将包含 AsyncScheduler 的根本修复，届时需要移除本 PR 添加的临时工作区。