

PR #6595 完整报告

verl-project/verl

[data] fix: default image_patch_size to processor's real patch_size in filter_overlong_prompts

合并时间: 2026-06-05 10:54

原文链接: <http://prhub.com.cn/verl-project/verl/pull/6595>

执行摘要

- 一句话: 修复 filter_overlong_prompts 硬编码图像 patch_size 导致的崩溃
- 推荐动作: 值得一读: 这是一个典型的多模态数据处理中配置不一致导致的 bug 修复, 展现了如何在数据预处理与 rollout 路径间对齐配置。虽然改动很小, 但设计上考虑了回退逻辑和显式 None 的情况, 适合作为类似问题的修复范例。

功能与动机

修复 Issue #6592: filter_overlong_prompts 硬编码 image_patch_size=14, 与 rollout 路径 (agent_loop.py:253) 使用的真实 processor.image_processor.patch_size 不一致, 导致视觉 token 计数偏差, 使过长样本通过过滤并在 rollout 时崩溃。

实现拆解

1. 在 verl/utils/dataset/rl_dataset.py 的 RLHFDataset.__init__ 方法中, 将硬编码的默认值 config.get("image_patch_size", 14) 替换为动态计算: 首先尝试从 processor.image_processor.patch_size 获取真实值, 若不可用则回退到 14。
2. 为避免显式设置为 None 时仍返回 None, 使用 config.get("image_patch_size") or _default_patch_size 确保任何假值 (包括 None) 都触发回退。
3. 无其他文件修改, 且不改变 API。

关键文件:

- verl/utils/dataset/rl_dataset.py (模块 数据集; 类别 source; 类型 core-logic; 符号 RLHFDataset.init): 核心修复文件: 修改 image_patch_size 的默认值获取方式, 从硬编码 14 改为动态获取 processor 的真实值。

关键符号: RLHFDataset.init

关键源码片段

verl/utils/dataset/rl_dataset.py

核心修复文件: 修改 image_patch_size 的默认值获取方式, 从硬编码 14 改为动态获取 processor 的真实值。

```
# 来自 verl/utils/dataset/rl_dataset.py 第 111-113 行
```

```
# 默认 patch_size 优先从 processor 获取真实值, 避免与 rollout 路径不一致
```

```
# 若 processor 或 image_processor 不存在，则回退 14
_default_patch_size = getattr(
    getattr(self.processor, "image_processor", None),
    "patch_size",
    14
)
# 使用 `or` 确保即使配置中显式设为 None 也能正确回退
self.image_patch_size = config.get("image_patch_size") or _default_patch_size
```

评论区精华

Gemini Code Assist 机器人建议使用 `config.get("image_patch_size") or _default_patch_size` 替代 `config.get("image_patch_size", _default_patch_size)`，以防止配置中显式设置 `null` 时返回 `None`。该建议被采纳并在第二次提交中实现。审核者 wuxibin89 批准了该 PR。

- 默认值处理避免 `None` 类型错误 (correctness): 采纳建议，第二次提交中改为 `config.get("image_patch_size") or _default_patch_size`。

风险与影响

- 风险：低风险。仅修改 RLHFDataset 中 `image_patch_size` 的默认值获取方式，且显式配置的 `image_patch_size` 仍优先。但需注意：若 `processor` 为 `None`，`getattr(self.processor, "image_processor", None)` 返回 `None`，继而 `getattr(None, "patch_size", 14)` 返回 14，行为与原硬编码一致。目前没有直接测试覆盖此逻辑，但已在 Qwen3-VL 上验证。
- 影响：影响范围有限：仅影响使用 RLHFDataset 且 `filter_overlong_prompts=True` 的视觉语言模型 (VLM) 训练流程，且模型 `image_processor.patch_size` 不等于 14 时。修复后，这些配置下的训练不再因过长短暂崩溃。
- 风险标记：缺少测试覆盖

关联脉络

- PR #6522 [vllm] fix: reset all caches after weight updates: 同为 vLLM rollout 路径的 bugfix，涉及多模态缓存问题，与此 PR 的 rollout 路径相关。
- PR #6554 [ci] chore: continue to replace the qwen25 model with the qwen3 model: 涉及 Qwen3-VL 模型测试，与此 PR 修复的 Qwen3-VL 场景相关。