

PR #7372 完整报告

verl-project/verl

[megatron] fix: Modify the VLM attention mask shape in TND format for NPU

合并时间: 2026-08-13 14:15

原文链接: <http://prhub.com.cn/verl-project/verl/pull/7372>

执行摘要

- 一句话: 修复 NPU 上 TND 格式 VLM attention mask 缺失
- 推荐动作: 值得精读: 对使用 NPU + Megatron + VLM (尤其是 Qwen3.5-VL) 的团队, 这是理解 `remove_padding` 下 VLM mask 数据契约的关键修复。值得关注的设计决策是在 mask 构造函数中统一保证非空, 而不是在每个 bridge 侧兜底, 属于最小侵入修复。建议后续跟进 `build_vlm_attn_mask_bshd` 的一致性处理与单元测试补充。

功能与动机

PR body 明确指出问题根因: 开启 `use_remove_padding` 且有 `vision_model` 的模型在 `mindspeed bridge` 下, GND 模块需要识别 padding attention mask; 若返回 `None`, padding 信息被丢弃, 后续构造 attention mask 时 padding 部分会被视为有效部分。VLM 需要在 `[B,S]` 格式的掩码中定位 `image` 等特殊标记, 因此必须移除 `build_vlm_attn_mask_thd` 中对 NPU 的提前返回, 恢复返回 `[B,S]` mask。

实现拆解

1. 修改 mask 构造函数 (`verl/models/mcore/util.py`): 在 `build_vlm_attn_mask_thd` 中删除 `if is_npu_available: return input_ids_rmpad, None` 分支, 所有后端统一通过 `offsets().diff()` 计算真实序列长度并构造 `[B,S]` 的 bool mask; 同时将 `input_ids_rmpad` 重命名为 `input_ids_with_pad`, 消除变量歧义。这样 Megatron/MindSpeed bridge 的 GND 模块能够拿到 padding 信息, VLM 也能在 `[B,S]` 掩码中定位 `image` 等特殊标记。
2. 更新 NPU 示例脚本 (`examples/grpo_trainer/run_qwen3_5_35b_megatron.sh`): 在 `npu` 分支追加 `actor_rollout_ref.actor.megatron.use_remove_padding=True`、`override_transformer_config.use_ascend_gdn=True`、`use_triton_gdn=False` 等配置, 并为 ACTOR/ROLLOUT/MODEL/REF 分别开启动态 bsz 与 `remove_padding` 开关; 同时移除 `set -xeuo pipefail`, 避免非关键命令失败直接退出。这些配置让修复在 Qwen3.5 35B Megatron NPU 场景中可复现。
3. 测试与配套: 本 PR 未新增单元或 e2e 测试, 修复正确性依赖 NPU 实机验证。建议后续补充针对 `build_vlm_attn_mask_thd` 的 CPU/GPU 单元测试 (构造嵌套张量并断言 mask 形状与取值), 并评估 `build_vlm_attn_mask_bshd` 是否也需要移除 NPU 分支以保持行为一致。

关键文件:

- verl/models/mcore/util.py (模块 模型工具; 类别 source; 类型 data-contract; 符号 build_vlm_attn_mask_thd) : 核心修复文件: 移除 NPU 提前返回 None 的分支, 使 TND 格式 VLM attention mask 在 NPU 上也以 [B,S] 形式返回, 避免 padding 信息丢失。
- examples/grpo_trainer/run_qwen3_5_35b_megatron.sh (模块 示例脚本; 类别 other; 类型 config) : 示例脚本: 为 NPU 用例补充 remove_padding、动态 bsz 与 Ascend GDN 配置, 使修复在 Qwen3.5 35B Megatron 场景中可复现。

关键符号: build_vlm_attn_mask_thd

关键源码片段

verl/models/mcore/util.py

核心修复文件: 移除 NPU 提前返回 None 的分支, 使 TND 格式 VLM attention mask 在 NPU 上也以 [B,S] 形式返回, 避免 padding 信息丢失。

```
# 修复后的 build_vlm_attn_mask_thd: 统一在 TND 路径上返回 [B,S] 掩码。
# NPU 上原先提前返回 (input_ids, None), 会使 Megatron / MindSpeed bridge
# 的 GND 模块丢失 padding 信息; 现在所有后端都走同一构造逻辑。
def build_vlm_attn_mask_thd(input_ids: torch.Tensor, pad_token_id: int = None):
    # 将嵌套 (jagged) 输入张量转为带 pad 的二维张量, pad_token_id 用于填充
    input_ids_with_pad = input_ids.to_padded_tensor(pad_token_id)
    # 每个样本的真实序列长度来自 offsets 的差分
    seqLens_in_batch = input_ids.offsets().diff()
    # 构造 bool 掩码: 真实 token 位置为 True, padding 位置为 False
    attention_mask = torch.zeros_like(input_ids_with_pad, dtype=torch.bool)
    for i, seqLen in enumerate(seqLens_in_batch):
        attention_mask[i, :seqLen] = True
    # 返回带 pad 的 token 序列与 [B,S] 掩码, 供 VLM 识别 image 等特殊标记
    return input_ids_with_pad, attention_mask
```

评论区精华

本 PR 未产生实质性 review 讨论; 唯一评论来自 CLAassistant 的 CLA 检查提醒 (已签署)。审查者 wucong25 直接批准, 无额外意见。PR body 中对设计权衡有明确说明: 返回 **None** 会让 padding 信息在 Megatron/MindSpeed bridge 的 GND 模块中被丢弃, 因此必须在 TND 路径上恢复 [B,S] mask。

- 暂无高价值评论线程

风险与影响

- 风险: 改动集中在 build_vlm_attn_mask_thd, 该函数被 TND 格式的 VLM 前向路径复用。风险包括: 1) build_vlm_attn_mask_bshd 仍保留 is_npu_available 提前返回, 与 TND 行为不一致, 未来可能复现同类问题; 2) 无测试覆盖, NPU + Megatron/MindSpeed 路径若依赖旧行为 (mask 为 None) 可能产生兼容性问题; 3) 示例脚本移除 set -xeuo pipefail, 配置错误时可能以非零码静默继续, 降低可诊断性; 4) NPU 上新增 mask 构造有微小计算开销, 但通常低于后处理收益。

- 影响：影响面主要是 NPU 上 Megatron/MindSpeed bridge 的 VLM（如 Qwen3.5-VL）在启用 `remove_padding` 后的 TND 训练 / 推理路径，修复了 padding 被误当有效 token 的静默错误；示例脚本让 Qwen3.5 35B Megatron NPU 用例可复现。影响程度中等偏小，不改变 API 和配置 schema，但改变了 NPU 上 mask 的输入形态，需要实机回归。对团队而言，本次改动小、易审阅，但缺少测试支撑，后续需补齐自动化覆盖。
- 风险标记：缺少测试覆盖，TND/BSHD 行为不一致，示例脚本移除 `set -e`，NPU 专用路径变更

关联脉络

- PR #7358 [megatron] feat: bucket packed sequence lengths: 同样修改了 `verl/models/mcore/util.py`，聚焦 Megatron 序列打包与掩码相关契约，与本 PR 的 TND mask 修复有文件交集。
- PR #7340 [megatron] fix: bugfix qwen 3 qwen 3.5 router replay: 同属 Qwen3.5 Megatron 修复线，本 PR 针对 Qwen3.5-VL 示例，说明该模型系列的 Megatron 路径正在持续加固。