

PR #7328 完整报告

verl-project/verl

[megatron] fix: forward mhc_multistream to MTP and skip activation reclaim for MTP checkpoints

合并时间: 2026-08-10 10:24

原文链接: <http://prhub.com.cn/verl-project/verl/pull/7328>

执行摘要

- 一句话: 修复 mHC+MTP 训练崩溃, 转发 multistream 并跳过 MTP 激活回收
- 推荐动作: 值得精读, 尤其是对 Megatron 显存管理与重计算机制感兴趣的同学。三个看点: 一是对 use-after-free 的定位思路——mhc_multistream 是 decoder 输出的别名、MTP depth 输入是其 torch.chunk view、MTP 与 decoder 的 backward 逆拓扑执行顺序, 三者叠加导致 `resize_(0)` 变成高危操作; 二是用闭包探测识别 checkpoint 层的技巧, 作者明确说明了 `__qualname__` 不带类名这一限制; 三是 `disable_mtp_completeness_check` 的 `nullcontext` 兜底模式, 是处理跨版本 API 差异的简洁写法。建议后续补充 CPU 层回归测试 (构造共享存储 view 场景) 以固化此修复, 并量化 MTP checkpoint 跳过回收后的显存影响。

功能与动机

PR body 明确指出此变更解决两个阻塞崩溃: 其一为前向 reshape 失败 (`RuntimeError: shape '[23986, 1, 4, 4096]' is invalid for input of size 98246656`), 根因是 verl 的 `_megatron_gptmodel_postprocess` 接受 `mhc_multistream` 参数却未转发给 `self.mtp(...)`, 使 MTP 内部收到 `None` 后走 `contracted` 路径, 与 mHC 分支期待的 `[s, b, n*h]` 张量形状不匹配; 其二为反向阶段的异步 CUDA illegal-memory-access, 根因是 MTP checkpoint 保存的 `hidden_states` 是 decoder `mhc_multistream` 的 `torch.chunk` view (共享存储), 而 MTP backward 按逆拓扑序先于 decoder 的 `learned_output_contract` backward 执行, `resize_(0)` 提前截断了 decoder 反向仍要读取的存储。

实现拆解

本变更按四个步骤拆解:

1. 前向修复 (`verl/models/mcore/mtp_patch.py`): `_megatron_gptmodel_postprocess` 新增 `mhc_multistream=None` 关键字参数, 并在构造 `mtp_kwargs` 时增加 `if mhc_multistream is not None: mtp_kwargs["mhc_multistream"] = mhc_multistream`, 使 MTP 模块能拿到收缩前的 `[s, b, n*h]` 多流张量作为 depth 输入, 与上游 `GPTModel._postprocess` 行为对齐。这是数据契约层面的补齐, MTP 内部分支逻辑本身无需改动。
2. 反向修复 (`verl/models/mcore/patch.py`): `patch_backward` (即 `rd.CheckpointFunction.backward` 的全局替换体) 在计算完梯度后、执行

`untyped_storage().resize_(0)` 激活回收前，先遍历 `ctx.run_function.__closure__` 的闭包单元，检查捕获对象是否为 `MultiTokenPredictionLayer` 实例；若是则跳过整段回收逻辑。原因是该 checkpoint 的 saved `hidden_states` 与 decoder 输出共享存储，而 MTP backward 先于 decoder backward 运行，截断会导致 decoder 反向读到已释放内存。非 MTP checkpoint 保持原有回收，保留 commit 04df110c 的 MoE 残差内存泄漏修复。

3. 配套兼容 (`verl/utils/vllm/vllm_quant_utils.py`)：两处改动。其一是 `FusedMoE` 导入从硬依赖改为 `try/except` 兜底置 `None`，因为 vLLM 0.26.1 彻底移除了 `FusedMoE` 名称（0.24.0 时已从类变为工厂函数），同时更新 `_MOE_STOP_CLASSES` / `_EXPERT_WEIGHT_CLASSES` 的分支判断注释，使其在 `FusedMoE` is `None` 时仍能通过 `RoutedExperts` / `MoERunner` 的兜底路径工作；其二是 `load_quantized_weights` 中调用 `model.load_weights(weights_quantized)` 前，用 `disable_mtp_completeness_check()` 上下文管理器包裹（通过 `try/except ImportError` 引入，旧版 vLLM 回退到 `nullcontext`），因为 RL refit 的权重是按桶分批到达的，完整性检查假设单次完整 checkpoint 加载会误报。
4. 测试与配置配套：本 PR 未新增或修改任何测试文件，也没有配置或部署改动。PR body 说明验证方式是针对 DeepSeek-V4 的 mHC + MTP 训练实验 (`use_fused_mhc=False + mtp.enable=True`)，未附显存基准或训练曲线。

关键文件：

- `verl/models/mcore/patch.py`（模块 模型层；类别 source；类型 core-logic；符号 `patch_backward`, `apply_patch_megatron_recomputation_backward`）：核心反向修复所在：全局 `CheckpointFunction.backward` 补丁中新增 MTP checkpoint 检测并跳过 `resize_(0)` 激活回收，是本次变更风险最高、机制最巧妙的部分。
- `verl/models/mcore/mtp_patch.py`（模块 模型层；类别 source；类型 data-contract；符号 `_megatron_gptmodel_postprocess`, `patch_postprocess`）：前向修复所在：`_megatron_gptmodel_postprocess` 新增 `mhc_multistream` 参数并转发给 MTP，解决 MTP 内部 shape 分支错误。
- `verl/utils/vllm/vllm_quant_utils.py`（模块 量化工具；类别 source；类型 dependency-wiring；符号 `load_quantized_weights`, `_MOE_STOP_CLASSES`, `_EXPERT_WEIGHT_CLASSES`）：配套兼容修复：`FusedMoE` 导入兜底适配 vLLM 0.26.1 移除类名，并禁用 RL refit 分桶加载下的 MTP 完整性检查。

关键符号：`patch_backward`, `_megatron_gptmodel_postprocess`, `load_quantized_weights`

关键源码片段

`verl/models/mcore/patch.py`

核心反向修复所在：全局 `CheckpointFunction.backward` 补丁中新增 MTP checkpoint 检测并跳过 `resize_(0)` 激活回收，是本次变更风险最高、机制最巧妙的部分。

```
# verl/models/mcore/patch.py | patch_backward 的后半段（梯度计算完成后）
cur_stream = torch.cuda.current_stream()
# 原有的激活内存回收逻辑（MoE 残差内存泄漏修复，commit 04df110c）：
# 对 checkpoint 的每个输入调用 untyped_storage().resize_(0) 提前释放存储。
# 但对 MTP 层 checkpoint 必须跳过：其保存的 hidden_states 是 decoder 输出的
# torch.chunk view (via make_viewless_tensor -> _kernel_make_viewless_tensor 的
```

```

# out.data = inp.data) , 与 mhc_multistream 共享同一块存储; 且 MTP backward
# 按逆拓扑序先于 decoder 的 learned_output_contract backward 执行, 此时
# resize_(0) 会截断 decoder 反向仍要读取的存储 -> 异步 CUDA illegal-memory-access。
is_mtp_checkpoint = False
run_fn = getattr(ctx, "run_function", None)
# 闭包探测: checkpoint_forward 的 custom_forward 闭包捕获了 MTP 层实例,
# 其 __qualname__ 是 ..._checkpointed_forward.<locals>.custom_forward, 不含类名,
# 所以只能遍历闭包单元检查捕获对象类型。
for cell in getattr(run_fn, "__closure__", None) or ():
    try:
        obj = cell.cell_contents
    except ValueError:
        continue
    if obj.__class__.__name__ == "MultiTokenPredictionLayer":
        is_mtp_checkpoint = True
        break
# 非 MTP checkpoint 保持原有回收, 避免破坏 MoE 泄漏修复;
# MTP checkpoint 则保留存储给 decoder 反向使用。
if not is_mtp_checkpoint:
    for t in detached_inputs:
        if isinstance(t, torch.Tensor) and t.requires_grad:
            t.record_stream(cur_stream)
            t.untyped_storage().resize_(0)
            if t.grad is not None:
                t.grad.record_stream(cur_stream)
                t.grad.untyped_storage().resize_(0)
# ctx.saved_tensors = None
return (None, None) + grads

```

verl/models/mcore/mtp_patch.py

前向修复所在: `_megatron_gptmodel_postprocess` 新增 `mhc_multistream` 参数并转发给 MTP, 解决 MTP 内部 shape 分支错误。

```

# verl/models/mcore/mtp_patch.py | _megatron_gptmodel_postprocess 的 MTP 调用部分
# 该补丁复制自上游 Megatron 的 GPTModel._postprocess, 用于支持 MTP、1f1b overlap 等特性。
# 本次新增 mhc_multistream 参数: mHC + MTP 时 decoder 返回
# (contracted_hidden, mhc_multistream), 后者是收缩前的 [s, b, n*h] 多流张量。
def _megatron_gptmodel_postprocess(
    self,
    hidden_states,
    input_ids,
    position_ids,
    labels,
    # ... 中间参数省略 ...
    is_spec_decode=None,
    mhc_multistream=None, # 新增: mHC 多流张量, 默认 None 保持旧行为
):
    ...
    if mtp_in_postprocess and labels is not None:

```

```

mtp_kwargs = dict(extra_block_kwargs or {})
if not hasattr(self.mtp, "_forward_has_padding_mask"):
    self.mtp._forward_has_padding_mask = "padding_mask" in signature(
        self.mtp.forward
    ).parameters
if self.mtp._forward_has_padding_mask:
    mtp_kwargs["padding_mask"] = padding_mask
# 关键修复：此前该参数被丢弃，MTP 内部收到 None 后走 contracted 分支，
# 与 mHC 分支的 _concat_embeddings 期待的 [s, b, n*h] 不匹配，触发 shape 崩溃。
if mhc_multistream is not None:
    mtp_kwargs["mhc_multistream"] = mhc_multistream
hidden_states = self.mtp(
    input_ids=input_ids,
    position_ids=position_ids,
    hidden_states=hidden_states,
    attention_mask=attention_mask,
    inference_params=inference_params,
    rotary_pos_emb=rotary_pos_emb,
    rotary_pos_cos=rotary_pos_cos,
    rotary_pos_sin=rotary_pos_sin,
    packed_seq_params=packed_seq_params,
    sequence_len_offset=sequence_len_offset,
    embedding=self.embedding,
    **mtp_kwargs,
)
...

```

verl/utils/vllm/vllm_quant_utils.py

配套兼容修复：FusedMoE 导入兜底适配 vLLM 0.26.1 移除类名，并禁用 RL refit 分桶加载下的 MTP 完整性检查。

```

# verl/utils/vllm/vllm_quant_utils.py
# FusedMoE 的兼容导入：vLLM 0.24.0 后 FusedMoE 从 nn.Module 变成工厂函数，
# vLLM 0.26.1 彻底移除了该名称。统一置 None 再走分支判断，避免 ImportError 拖垮整个模块。
try:
    from vllm.model_executor.layers.fused_moe.layer import FusedMoE
except ImportError:
    FusedMoE = None

```

```

def load_quantized_weights(weights, model_runner, is_drafter=False):

```

```

    ...
    # 最终把权重加载进 vLLM。
    # MTP 完整性检查 (disable_mtp_completeness_check) 假定权重来自单次完整
    # checkpoint 加载；而 RL refit 场景下权重是按桶分批到达的，会误报缺失，
    # 因此需要禁用——这与 vLLM 自身 NCCL/IPC 引擎的做法一致。
    # nullcontext 兜底覆盖缺少该 API 的旧版 vLLM。
    try:
        from vllm.model_executor.model_loader.mtp_validation import (

```

```

        disable_mtp_completeness_check,
    )
except ImportError:
    disable_mtp_completeness_check = nullcontext
try:
    with disable_mtp_completeness_check():
        loaded_params = model.load_weights(weights_quantized)
finally:
    # 还原参数类型伪装 (subclass_type)
    for name, param in model.named_parameters():
        if hasattr(param, "orig_type"):
            param.__class__ = param.orig_type
            del param.orig_type
    return loaded_params

```

评论区精华

本 PR 没有任何实质性的 review 讨论线程：Copilot reviewer 因配额限制未能生成评审（"Copilot was unable to review... quota limit"），维护者 wuxibin89 直接批准（APPROVED），无 review 评论、无 issue 评论。技术细节与设计权衡全部沉淀在 PR body 中，包括对两个崩溃根因的深入分析、上游 Megatron 行为对照，以及对共享存储与反向拓扑顺序的论证。

- 暂无高价值评论线程

风险与影响

- 风险：具体风险如下：

1. 闭包探测的脆弱性（`verl/models/mcore/patch.py`）：通过 `obj.__class__.__name__ == "MultiTokenPredictionLayer"` 字符串比较识别 MTP checkpoint，依赖 Megatron 内部类名稳定。若上游改名或 MTP 实现重构，检测会静默失效——失效方向一是重新引入异步 CUDA illegal-memory-access（崩溃），二是误放行其他 checkpoint 导致 MoE 残差内存泄漏回归（commit 04df110c 的修复被绕过）。此外 `__closure__` 是 CPython 实现细节，`cell.cell_contents` 访问有 `ValueError` 兜底，但整体机制比较 hacky。
2. 内存回收语义变化：跳过 MTP checkpoint 的 `resize_(0)` 后，MTP 层激活内存不再被及时截断释放，长序列或大 batch 下峰值显存可能上升。PR body 未提供显存前后对比数据，无法量化影响。
3. 全局补丁的波及面：`patch.py` 中 `rd.CheckpointFunction.backward = patch_backward` 是模块级全局替换，影响所有 Megatron 重计算 checkpoint 的 backward。虽然行为仅在 MTP checkpoint 时变化，但本 PR 修改的热路径每次 backward 都新增了闭包遍历的开销（极小，但有）。
4. MTP 完整性检查被禁用（`vllm_quant_utils.py`）：RL refit 分桶加载确实无法通过完整性检查，禁用是合理的；但这也意味着加载缺失会被推迟到运行时才暴露。风险通过 `nullcontext` 兜底和注释说明做了控制，旧版本 vLLM 不受影响。
5. 缺少测试覆盖：本 PR 没有配套单测或 e2e 测试。这类共享存储 + 视图别名问题是可以 CPU 单测复现的（如构造 `torch.chunk` view 后执行 `resize_(0)` 验证后续读取），当前只能

依赖实验验证，回归防护不足。 - 影响：影响范围评估：

- 直接受益用户：在 Megatron-core 上以 `use_fused_mhc=False + mtp.enable=True` 训练 DeepSeek-V4 的团队，此前 mHC + MTP 组合完全不可用（前向崩溃或反向异步崩溃），本 PR 解除该阻塞。
- 间接影响面：patch.py 的重计算 backward 补丁是所有 Megatron 训练共用的全局路径，改动虽小但在核心热路径上；mtp_patch.py 的 `_postprocess` 补丁影响所有启用 MTP 的 Megatron 模型（即使无 mHC，新增参数为 None 时行为不变）。
- vLLM refit 路径：vllm_quant_utils.py 影响所有走量化权重加载的 rollout refit 流程，改动为向后兼容设计，风险低。
- 团队协作：HollowMan6 单人提交，2 个 commit（主体修复 + 新版本兼容），评审由 wuxibin89 一人批准，无多轮博弈。
- 风险标记：核心路径变更，缺少测试覆盖，依赖字符串匹配探测模块类型，内存回收策略受执行顺序影响

关联脉络

- PR #7241 [megatron, hardware] fix: pure-torch fast_hadamard_transform fallback for DSA on ROCm: 同样修改 verl/models/mcore/patch.py，属于同一 Megatron 补丁文件上的连续修复线。
- PR #7224 [vllm] feat: enhance DeepSeek V4 fp8/fp4 linear and moe weight refit: 同改 verl/utils/vllm/vllm_quant_utils.py 及 DeepSeek V4 权重 refit 路径，本 PR 的 MTP 完整性检查禁用是对该 refit 能力的补充。
- PR #7221 [megatron] feat: support contiguous context-parallel layout for DeepSeek V4: 同为 DeepSeek V4 在 Megatron 上的支持线（涉及 verl/models/mcore/util.py 与 transformer_impl.py），本 PR 的 mHC + MTP 修复延续这一功能演进方向。