

PR #5599 完整报告

verl-project/verl

[megatron] fix: Qwen3.5 LoRA & MTP support (with Megatron-Bridge)

合并时间: 2026-07-21 14:03

原文链接: <http://prhub.com.cn/verl-project/verl/pull/5599>

执行摘要

- 一句话: 修复 Qwen3.5 LoRA 与 MTP 权重同步与分桶传输问题
- 推荐动作: 值得精读, 尤其是权重名称归一化的重构思路 (将名称处理推向接收端) 和多桶 LoRA 累积的设计; 对于维护后续模型兼容性有参考价值; 需关注 Copilot 提出的性能优化建议 (reverse_mapping 缓存) 是否已在最终版本完全采纳。

功能与动机

之前权重同步路径依赖发送端名称重写, 硬编码 `.base_layer` 处理, 对 packed projections 和 fused MoE 模型脆弱; 异步 LoRA 更新假设每个 adapter 在一个 IPC bucket 中到达, 但分桶传输不保证, 可能导致 `add_lora` 请求不完整; 同时也存在 Qwen 嵌套 `text_config` 的 MTP 字段、新版 Megatron-LM `process_mtp_loss` API 等兼容性问题。

实现拆解

1. 接收端权重名称归一化 (verl/workers/rollout/vllm_rollout/utils.py) : 新增 `_resolve_weight_name_for_vllm`、`_iter_packed_owner_weight_names`、`_strip_bridge_base_layer_from_expert_alias` 等静态方法, 在 vLLM 端动态解析传入的权重名称, 通过 `resolve_base_layer_name` 探测目标命名空间, 替代发送端硬编码。
2. 多桶 LoRA 更新支持 (verl/workers/rollout/vllm_rollout/bucketed_weight_transfer.py) : `BucketedWeightReceiver.iter_weights` 新增流式生成器, 保留分桶背压语义; `update_weights_from_ipc` 中累积跨桶 LoRA 张量并在最终桶到达后调用 `add_lora`, 避免并发混叠。
3. PEFT 工具重构 (verl/utils/megatron_peft_utils.py) : 移除硬编码的 `STACKED_PARAMS` 列表, 替换为 `add_base_layer_to_name` / `remove_base_layer_from_name` / `resolve_base_layer_name` 通用函数; 补充 GDN 模块映射 (`in_proj -> in_proj_qkv` 等)。
4. MTP 兼容性更新 (verl/workers/megatron_workers.py、verl/utils/megatron_utils.py) : MTP 检测同时支持 `num_nextn_predict_layers` 和 `mtp_num_hidden_layers`, 并能从嵌套 `text_config` 中读取; `check_mtp_config` 文档同步更新。
5. Worker/runtime 修复 (verl/workers/engine/megatron/transformer_impl.py、verl/trainer/ppo/v1/trainer_base.py) : 仅在 param offload 启用时才将 actor engine 卸载到 CPU; `vision_config` 在 checkpoint 备份中保持; 修正 missing speculative

decoding metrics 时的安全处理。

关键文件：

- verl/workers/rollout/vllm_rollout/utils.py (模块 权重同步；类别 source；类型 core-logic；符号 `_map_weight_name_for_vllm`, `_is_leaf_weight_or_bias_name`, `_strip_bridge_base_layer_from_expert_alias`, `_adapt_weight_names_for_model`)：核心变更文件：实现接收端权重名称归一化，新增 8 个关键方法替换硬编码名称重写。
- verl/utils/megatron_peft_utils.py (模块 PEFT 工具；类别 source；类型 core-logic；符号 `add_base_layer_to_name`, `remove_base_layer_from_name`, `resolve_base_layer_name`)：重构了 `base_layer` 处理，使用通用函数替代硬编码参数列表，新增 GDN 映射。
- tests/utils/test_vllm_weight_name_normalization_on_cpu.py (模块 名称映射测试；类别 test；类型 test-coverage；符号 `_FakeMapper`, `_FakeModel`, `_make_worker`, `test_normalize_base_sync_weight_names_preserves_expert_logical_aliases`)：新增测试覆盖核心权重名称归一化逻辑，确保 MoE expert 别名和 packed module 映射正确。
- tests/utils/test_bucketed_weight_transfer.py (模块 分桶传输测试；类别 test；类型 test-coverage；符号 `test_iter_weights_retains_direct_tensor_until_ack`, `_FakeDevice`, `_FakeSocket`)：补充 `iter_weights` 保留直接张量直到确认的回归测试，覆盖 IPC handle 路径和垃圾回收时序。
- verl/workers/rollout/vllm_rollout/bucketed_weight_transfer.py (模块 分桶传输；类别 source；类型 core-logic；符号 `iter_weights`)：新增 `iter_weights` 流式接收方法，支持层状重载和跨桶 LoRA 累积。
- verl/workers/engine/megatron/transformer_impl.py (模块 引擎适配；类别 source；类型 dependency-wiring)：移除了对 `add_base_layer_suffix` 的调用，将归一化逻辑完全转移到接收端；修正 LoRA 合并时的 offload 条件。

关键符号：`_map_weight_name_for_vllm`, `_is_leaf_weight_or_bias_name`, `_strip_bridge_base_layer_from_expert_alias`, `_adapt_weight_names_for_model`, `_iter_packed_owner_weight_names`, `_resolve_weight_name_for_vllm`, `_candidate_exists`, `_iter_normalized_base_sync_weights`, `add_base_layer_to_name`, `remove_base_layer_from_name`, `resolve_base_layer_name`, `iter_weights`

关键源码片段

verl/workers/rollout/vllm_rollout/utils.py

核心变更文件：实现接收端权重名称归一化，新增 8 个关键方法替换硬编码名称重写。

```
@staticmethod
def _resolve_weight_name_for_vllm(model, weight_name: str):
    """Resolve a weight name by trying the vLLM mapper, then probing the namespace."""
    # Try vLLM's built-in weight name mapper (e.g. for HF -> vLLM conversion)
    mapped = vLLMColocateWorkerExtension._map_weight_name_for_vllm(model, weight_name)
    if mapped is not None and mapped != weight_name:
        return mapped

    # If we have a packed owner like qkv_proj, we need to map to the packed name.
```

```

for candidate in vLLMColocateWorkerExtension._iter_packed_owner_weight_names(model,
weight_name):
    if vLLMColocateWorkerExtension._candidate_exists(model, candidate):
        return candidate

# Last resort: probe the live model namespace.
if not vLLMColocateWorkerExtension._candidate_exists(model, weight_name):
    # If the name is a leaf weight/bias that doesn't exist, try adding/removing .base_layer
    if vLLMColocateWorkerExtension._is_leaf_weight_or_bias_name(weight_name):
        from verl.utils.megatron_peft_utils import resolve_base_layer_name
        return resolve_base_layer_name(weight_name,
            exists=lambda n: vLLMColocateWorkerExtension._candidate_exists(model, n))
return weight_name

@staticmethod
def _iter_packed_owner_weight_names(model, weight_name: str):
    """Yield packed-owner names for unpacked HF aliases such as q/k/v proj."""
    packed_modules_mapping = getattr(model, "packed_modules_mapping", None) or {}
    if not packed_modules_mapping or "." not in weight_name:
        return

    reverse_mapping: dict = {}
    for packed_name, unpacked_names in packed_modules_mapping.items():
        for unpacked_name in unpacked_names:
            reverse_mapping.setdefault(unpacked_name, []).append(packed_name)

    parts = weight_name.split(".")
    module_idx = -3 if len(parts) >= 3 and parts[-2] == "base_layer" else -2
    leaf = parts[-1]
    module_name = parts[module_idx]
    if leaf not in {"weight", "bias"}:
        return

    for packed_name in reverse_mapping.get(module_name, []):
        parts[module_idx] = packed_name
        yield ".".join(parts)

```

tests/utils/test_bucketed_weight_transfer.py

补充 `iter_weights` 保留直接张量直到确认的回归测试，覆盖 IPC handle 路径和垃圾回收时序。

```

def test_iter_weights_retains_direct_tensor_until_ack(monkeypatch):
    # ... ( setup: fake socket with direct-ipc handle )

    receiver = BucketedWeightReceiver("ipc:///tmp/unused.sock", device=torch.device("cpu"), use_
shm=False)
    fake_socket = _FakeSocket()
    monkeypatch.setattr(receiver, "_init_socket", lambda: setattr(receiver, "socket", fake_socket))
    monkeypatch.setattr(receiver, "_init_buffer", lambda: setattr(receiver, "buffer", torch.empty(0)
))

```

```

monkeypatch.setattr(bucketed_weight_transfer, "rebuild_ipc", _rebuild_ipc)
monkeypatch.setattr(bucketed_weight_transfer, "get_torch_device", lambda: _FakeDevice)

weights = receiver.iter_weights()
name, tensor = next(weights)
assert name == "large.weight"
torch.testing.assert_close(tensor, torch.arange(6, dtype=torch.float32).view(2, 3))
del tensor

with pytest.raises(StopIteration):
    next(weights)

assert acknowledged # ACK 在张量引用释放后才发送

```

评论区精华

- `_is_leaf_weight_or_bias_name` 对 MoE expert 权重判定过严: Copilot 指出该函数将 `w13_weight` 等视为非 leaf, 导致 `_strip_bridge_base_layer_from_expert_alias` 错误剥离 `base_layer`。作者已修复为匹配 `_weight/_bias` 后缀。
- `_iter_packed_owner_weight_names` 每次重建 `reverse_mapping` 影响性能: Copilot 建议缓存, 作者响应 fixed。
- `iter_weights` 忽略直接发送大权重 IPC handle: Copilot 指出 `iter_weights` 未处理 `handle is not None` 分支导致大权重丢失。作者修复。
- 跨桶 LoRA 张量必须 clone: Gemini Code Assist 指出 `lora_weights` 中的张量是共享 buffer 的视图, 后续 bucket 写入会覆盖。作者添加 `.clone()`。
- 冗余 `assert self.device is not None`: Gemini 建议删除, 因为后续 `rebuild_ipc` 会自然失败。作者移除。
- MTP 断言消息不明确: Copilot 指出消息未提及 `mtp_num_hidden_layers`。作者更新。
- `check_mtp_config` 文档字符串未更新: 作者 fixed。
 - `_is_leaf_weight_or_bias_name` 对 MoE expert 权重的判定 (correctness): 作者已将条件扩展为识别 `_weight/_bias` 后缀。
 - `_iter_packed_owner_weight_names` 每次重建 `reverse_mapping` 性能问题 (performance): 作者回应 fixed (具体解决方案可能为缓存或预计算)。
 - `iter_weights` 忽略直接发送大权重的 IPC handle (correctness): 作者修复, 添加 `handle is not None` 分支调用 `rebuild_ipc`。
 - 跨桶 LoRA 张量需 clone 避免数据损坏 (correctness): 作者在 commit fb75752 中添加 `.clone()`。
 - 冗余 `assert self.device is not None` (style): 作者移除断言。
 - MTP 断言消息未提及 `mtp_num_hidden_layers` (documentation): 作者更新消息。

风险与影响

- 风险:

1. LoRA 数据损坏风险：多桶累积的 LoRA 张量若不 clone 会被后续 bucket 覆盖，已在 commit fb75752 修复。
2. 分桶传输正确性：iter_weights 新路径对大权重 IPC handle 处理不完整曾引入 bug (commit f278933d3c38 修复)。
3. 性能退化：_iter_packed_owner_weight_names 每次调用重建 reverse_mapping，虽已加 TODO，但在大批量权重同步时可能成为热点。
4. 测试覆盖局限：测试主要在 CPU 上运行，未覆盖真实多 GPU 分布式场景，分桶传输的竞争条件可能漏测。
5. MTP 配置解析风险：嵌套 text_config 的 fallback 逻辑可能在某些自定义配置中产生误判，需进一步用户反馈。- 影响：影响 Qwen3.5 及其变体在 Megatron 后端下的 LoRA 微调和 MTP 推理；修复了异步 LoRA 更新中可能的数据损坏，提升稳定性；接收端名称归一化使未来模型适配更灵活，但依赖 Megatron-Bridge 和 vLLM 的上下游 PR 同步；同步效率无显著下降，但名称解析计算略有增加。- 风险标记：核心路径变更，多桶依赖，缺少分布式测试，性能回归风险

关联脉络

- PR #7014 [fsdp] fix: sync merged LoRA weights before context exit: 同为 LoRA 权重同步的 bugfix，但作用于 fsdp 后端；本 PR 针对 Megatron 后端，两者共同提升了 LoRA 场景的可靠性。
- PR #6660 [model, fsdp] fix: support Qwen3.5 linear attention under Ulysses SP: 同属 Qwen3.5 系列支持，本 PR 延伸了 LoRA 和 MTP 支持，两者协作完善 Qwen3.5 在 verl 中的功能。
- PR #6974 [ckpt, rollout] feat: sharded delta weight sync over NCCL for disaggregated rollout: 涉及权重同步引擎改进，本 PR 在权重同步方面进行了类似的企业级增强（接收端名称归一化、多桶 LoRA 累积）。