

PR #7413 完整报告

verl-project/verl

[sglang] fix: lora sglang e2e

合并时间: 2026-08-14 17:49

原文链接: <http://prhub.com.cn/verl-project/verl/pull/7413>

执行摘要

- 一句话: 修复 LoRA 同步至 SGLang 的七处链式缺陷
- 推荐动作: 值得精读。PR body 是“链式缺陷排查”的范本: 逐项回滚的负向对照能区分“确证修复”与“防御修复”, `wake_up()` 一节如实承认没有证据证明必要但仍保留, 实证纪律非常值得借鉴。两个设计原则可以沉淀为团队约定: ①让不一致的调用点向已经能工作的那个对齐, 而不是发明第三种答案; ②对无法表达的配置 (正则 `target_modules`) 选择启动时响亮失败, 而不是静默适配错误模块集。重点阅读 `utils.py` 的三个 helper、`sglang_rollout.py` 的 `wrap_lora_params()`, 以及 `http_server_engine.py` 的错误 body 日志改造。

功能与动机

PR body 明确指出这一组合从未完整跑通: “Makes LoRA adapter sync to SGLang work end to end, on both the FSDP and the megatron path. This pairing had never been run to completion, so its defects sit in a chain: each one is only reachable once the previous is fixed, and a fix for any single layer looks like it changes nothing.” 而 `examples/tuning/lora/` 下 5 个 LoRA 示例全部用 vLLM、没有任何 SGLang 示例, 正是这条长链条长期藏匿的原因。此外本 PR 关闭 #7289 (adapter 模式 resume 了从未释放的 `weights` 标签导致 `KeyError: 'weights'`) 与 #7290 (megatron 的 `peft_config` 缺 `peft_type`)。

实现拆解

1. 统一配置入口与 `target_modules` 翻译 (缺陷 1/7)

`utils.py` 新增三个核心 helper: `normalize_peft_config_for_sglang()` 接受 dict (与 `BaseEngine.get_per_tensor_param` 声明一致), 展开 `task_type/peft_type` 枚举为字符串, 缺失 `peft_type` 时报错并列出已有键 (点名 #7290 的生产者), 且对字符串形态的 `target_modules` 原样保留——上一版无条件 `list()` 会再次把 'all-linear' 拆成字符。
`sglang_lora_target_modules()` 在启动参数侧把 'all-linear' 翻译为 SGLang 的 'all' 哨兵, 对正则字符串直接抛错 (fail-loud)。`async_sglang_server.py` 的 `launch_server()` 改用这两个 helper, 并用 `lora_rank_of()` 同时兼容 megatron 的 `model.lora.rank` 与 FSDP 的扁平 `model.lora_rank` (此前 megatron 下 `max_lora_rank` 恒为 0, SGLang 拒绝启动)。

2. 修复 adapter 载荷构造与 HTTP 线协议 (缺陷 2/3/7)

sglang_rollout.py 的 wrap_lora_params() 签名从 LoraConfig 改为 dict, 去掉 asdict(), 改走 normalize_peft_config_for_sglang(); 请求字段从 serialized_tensors 改为 serialized_named_tensors (每个 TP rank 一份副本), 与同文件基础权重同步保持同一形态, 对齐 `sglang >= 0.5.18` 的 List[bytes]。http_server_engine.py 的 load_lora_adapter_from_tensor() 与 update_weights_from_tensor() 一致地对每份 payload 做 base64 编码。

3. 修复 sleep/resume 状态机不对称 (缺陷 4)

sleep_level 改为构造期从 lora_served_as_adapter(model_config) 直接推导, 不再等第一次同步后才由 engine_workers.update_weights() 赋值, 保证 sleep、resume、engine_workers 三处读取点从第 0 步一致。engine_workers.py 在 sleep_level == 1 时跳过 resume(["weights"]), 并用 getattr(self.rollout, 'sleep_level', 2) 保证非 SGLang 后端行为不变。async_sglang_server.py 的 wake_up() COLOCATED 分支按 lora_as_adapter 选择 tags, 与 sleep() 对称 (作者明示这是防御性修复, 当前配置下不可达)。

4. 适配 FSDP 侧张量放置与命名 (缺陷 5/6)

新增 _to_ipc_device(): collect_lora_params() 刻意把张量放 CPU 以压低峰值显存, 但 SGLang 的 IPC 补丁只给设备张量登记 reducer 槽位, 序列化前需逐个搬回设备。新增 _strip_lora_base_layer(): 基础权重同步时去掉训练引擎为 vLLM 插入的 .base_layer 段。最后一笔 commit 还把硬编码 torch.cuda 换成 verl 的 get_device_id(), 通过平台层解析当前加速器, 避免 check_device_api_usage 拦截。

5. Megatron 契约补齐与验证策略

megatron_peft_utils.py 的 build_peft_config_for_vllm() 补上 peft_type: PeftType.LORA (#7290), 使同一个 dict 同时满足 vLLM 与 SGLang 两个消费方; engine/base.py 的 get_per_tensor_param() docstring 写明返回 dict 的必需键及枚举成员事实, 防止两个生产者再次漂移。测试配套说明: 不写单测 (线协议缺陷需要活体服务器, mock 测试等于测试 mock 的形状), 改用 Qwen3-0.6B/GSM8K 4 × A100 的 e2e 跑批加五项负向对照; 提交历史中还删除了从上游 #7288/#7287 带来的三个 CPU 单测文件 (commit c4c105c), 理由是行为已由 e2e 覆盖且断言属于上游 PR。

关键文件:

- verl/workers/rollout/sglang_rollout/sglang_rollout.py (模块 推理引擎; 类别 source; 类型 core-logic; 符号 _strip_lora_base_layer, _to_ipc_device, wrap_lora_params): 核心修复文件: wrap_lora_params() 从 asdict() 切换到 normalize_peft_config_for_sglang(), 请求字段改为 serialized_named_tensors 并与基础权重同步对齐; 新增 _strip_lora_base_layer() 与 _to_ipc_device() 分别处理 .base_layer 残留与 CPU 张量 IPC; sleep_level 改为构造期从配置推导, 闭环缺陷 4。
- verl/workers/rollout/sglang_rollout/utils.py (模块 推理引擎; 类别 source; 类型 core-logic; 符号 normalize_peft_config_for_sglang, lora_rank_of, sglang_lora_target_modules): 新增三个核心 helper: normalize_peft_config_for_sglang() (枚举展开 + peft_type 守卫 + 保留字符串 target_modules)、sglang_lora_target_modules() ('all-linear' → 'all' 翻译、正则报错)

、lora_rank_of()（兼容两套永不同步的 LoRA 配置块），是缺陷 1/7 与 megatron 启动参数侧修复的汇聚点。

- verl/workers/rollout/sglang_rollout/http_server_engine.py（模块 推理引擎；类别 source；类型 core-logic；符号 _log_error_body, load_lora_adapter_from_tensor）：
load_lora_adapter_from_tensor() 改用 serialized_named_tensors 并逐份 base64 编码（与 update_weights_from_tensor 对齐）；新增 _log_error_body() 记录失败响应 body，aiohttp 路径必须在 async with 内读完 body，否则缺陷 7 的 400 只呈现状态行、原因被丢弃。
- verl/workers/rollout/sglang_rollout/async_sglang_server.py（模块 推理引擎；类别 source；类型 dependency-wiring；符号 wake_up, launch_server）：launch_server() 的 LoRA 参数改用 lora_rank_of() 与 sglang_lora_target_modules()，修复 megatron 下 max_lora_rank 恒为 0 的启动失败；wake_up() 的 COLOCATED 分支按 lora_as_adapter 选择 tags，与 sleep() 对称（防御性修复，当前配置下不可达）。
- verl/utils/megatron_peft_utils.py（模块 适配器配置；类别 source；类型 dependency-wiring；符号 build_peft_config_for_vllm）：build_peft_config_for_vllm() 补上 peft_type: PeftType.LORA，修复 #7290：同一个 peft_config dict 要同时满足 vLLM 的 PEFTHelper 与 SGLang 的 adapter loader，而 SGLang loader 强制要求 peft_type 键。
- verl/workers/engine/base.py（模块 引擎基类；类别 source；类型 core-logic；符号 get_per_tensor_param）：把 get_per_tensor_param() 返回 dict 的契约（必需键、枚举成员而非字符串）写进 docstring，让任何序列化该 dict 的消费方知道要 unwrap 枚举，防止两个生产者再次漂移。
- verl/workers/engine_workers.py（模块 引擎编排；类别 source；类型 core-logic；符号 update_weights）：adapter 模式（sleep_level == 1）跳过 resume(["weights"]), 闭环缺陷 4 的调用侧：sleep() 只释放了 kv_cache，resume 却在请求从未释放的 weights 标签，SGLang 的 offload_tags.remove() 是严格集合操作直接 KeyError；getattr 默认 2 保证非 SGLang 后端行为不变。

关键符号：wrap_lora_params, normalize_peft_config_for_sglang,
sglang_lora_target_modules, lora_rank_of, _strip_lora_base_layer, _to_ipc_device,
_log_error_body, load_lora_adapter_from_tensor, build_peft_config_for_vllm, wake_up

关键源码片段

verl/workers/rollout/sglang_rollout/sglang_rollout.py

核心修复文件：wrap_lora_params() 从 asdict() 切换到 normalize_peft_config_for_sglang()，请求字段改为 serialized_named_tensors 并与基础权重同步对齐；新增 _strip_lora_base_layer() 与 _to_ipc_device() 分别处理 .base_layer 残留与 CPU 张量 IPC；sleep_level 改为构造期从配置推导，闭环缺陷 4。

```
# sleep_level 由构造期从配置推导，而不是等第一次同步后再赋值：
# sleep() 的释放分支与 resume() 的恢复分支从此共享同一个谓词（缺陷 4）。
self.sleep_level = 1 if lora_served_as_adapter(self.model_config) else 2

def _strip_lora_base_layer(name: str) -> str:
```

```
# 训练引擎为 vLLM 插入的 .base_layer 段，SGLang 没有这一层，不剥掉会 KeyError（缺陷 6）
return name.replace('.base_layer.', '.') if '.base_layer.' in name else name
```

```
def _to_ipc_device(tensor: torch.Tensor) -> torch.Tensor:
    # collect_lora_params() 刻意把张量放 CPU 以压低峰值显存，但 SGLang 的 IPC 补丁
    # 只给设备张量登记 reducer 槽位，所以序列化前要逐个搬回设备（缺陷 5）
    return tensor.to(get_device_id(), non_blocking=True) if tensor.device.type == 'cpu' else
    tensor
```

```
def wrap_lora_params(self, peft_config: dict, weights: Generator[tuple[str, torch.Tensor]]):
    # 两个引擎交给这里的都是 dict（FSDP 是 LoraConfig.to_dict()），不再调用
    # asdict()；枚举展开与 target_modules 保留统一由 helper 处理（缺陷 2/7）。
    peft_config_json = normalize_peft_config_for_sglang(peft_config)

    processed_weights: dict[str, torch.Tensor] = {
        name: _to_ipc_device(_preprocess_tensor_for_update_weights(tensor.detach())) for name,
        tensor in weights
    }

    # 每个 TP rank 一份序列化副本，字段与同文件基础权重同步一致：
    # sglang >= 0.5.18 的 serialized_named_tensors 是 List[bytes]（缺陷 3）。
    tp_size = self.device_mesh['infer_tp'].mesh.size()[0]
    serialized_named_tensors = [MultiprocessingSerializer.serialize(processed_weights) for _ in
    range(tp_size)]

    return peft_config_json, serialized_named_tensors
```

verl/workers/rollout/sglang_rollout/utils.py

新增三个核心 helper：normalize_peft_config_for_sglang()（枚举展开 + peft_type 守卫 + 保留字符串 target_modules）、sglang_lora_target_modules()（'all-linear' → 'all' 翻译、正则报错）、lora_rank_of()（兼容两套永不同步的 LoRA 配置块），是缺陷 1/7 与 megatron 启动参数侧修复的汇聚点。

```
def normalize_peft_config_for_sglang(peft_config: dict) -> dict:
    # 两个引擎返回的 task_type / peft_type 是枚举成员，跨 HTTP 前必须换成字符串
    normalized = dict(peft_config)
    for key in ('task_type', 'peft_type'):
        if key in normalized:
            normalized[key] = getattr(normalized[key], 'value', normalized[key])
    if 'peft_type' not in normalized:
        # megatron 引擎一度缺这个键（#7290），SGLang loader 强制要求，直接点名报错
        raise ValueError(
            'adapter config has no peft_type, which SGLang adapter loader requires; '
            'keys present: ' + ', '.join(sorted(normalized))
        )
    # 裸字符串必须原样保留：list() 会把 'all-linear' 拆成单个字符（缺陷 7 的形状）
    target_modules = normalized['target_modules']
```

```
normalized['target_modules'] = target_modules if isinstance(target_modules, str) else
list(target_modules)
return normalized
```

```
def lora_rank_of(model_config) -> int:
    # LoRA rank 可能落在两套从未同步的配置块里：megatron 用 model.lora.rank,
    # FSDP 用扁平 model.lora_rank, 取两者最大值兜底（megatron 启动参数侧修复）
    return max(int(getattr(model_config, 'lora_rank', 0) or 0), int(model_config.lora.get('rank', 0)
    or 0))
```

```
def sglang_lora_target_modules(target_modules: Any) -> list[str]:
    # PEFT 的 'all-linear' 简写在 SGLang 里写作 'all', 由服务端自行展开（缺陷 1）
    if target_modules == 'all-linear':
        return ['all']
    if isinstance(target_modules, str):
        # PEFT 把正则串对整个参数键做 fullmatch, SGLang 无法表达；
        # 与其静默适配到一组不同的模块，不如启动时响亮报错
        raise ValueError('SGLang cannot serve a regex target_modules; use all-linear or list
        module names.')
    return list(target_modules)
```

verl/workers/rollout/sglang_rollout/http_server_engine.py

load_lora_adapter_from_tensor() 改用 serialized_named_tensors 并逐份 base64 编码（与 update_weights_from_tensor 对齐）；新增 _log_error_body() 记录失败响应 body，aiohttp 路径必须在 async with 内读完 body，否则缺陷 7 的 400 只呈现状态行、原因被丢弃。

```
# 失败响应 body 的日志上限，防止刷屏
_ERROR_BODY_CHARS = 2000
```

```
def _log_error_body(endpoint: str, status: int, body: str) -> None:
    # raise_for_status() 抛出的错误只带状态行，真正原因（pydantic 校验报告）在 body 里
    logger.error(f'HTTP error for {endpoint}: {status}, body={body[:_ERROR_BODY_CHARS]}')
```

```
async def load_lora_adapter_from_tensor(self, req):
    import base64
    # 与 update_weights_from_tensor() 保持一致：字段是 List[bytes], JSON 没有 bytes,
    # 逐份 base64 编码（缺陷 3 的 HTTP 侧）；缺陷 7 的 400 凭 body 直接可读
    serialized_named_tensors = [base64.b64encode(t).decode('utf-8') for t in req.serialized_
    named_tensors]
    return await self._make_async_request(
        'load_lora_adapter_from_tensors',
        {
            'lora_name': req.lora_name,
            'config_dict': req.config_dict,
            'serialized_named_tensors': serialized_named_tensors,
```

```
},  
)
```

评论区精华

本 PR 没有任何 review 评论 (review_comments_count 为 0)，唯一的 issue 评论是 wuxibin89 的格式请求，已按 CONTRIBUTING 的 code-linting-and-formatting 处理，最终由 wuxibin89 APPROVED。实质性讨论全部写在 PR body 中，作者主动向 reviewer 交代了关键取舍：缺陷 3 是与 #7287 的唯一分歧点 (sglang 将字段从 serialized_tensors: str 演进到 serialized_named_tensors: List[bytes]，#7287 只对 0.5.8 正确)；wake_up() 对称性是实证无法复现的防御性修复；单测缺位的原因是“mock-level tests for this chain assert the shape of the mock”；.base_layer 剥离是权宜修复，引擎按后端分支“would be cleaner, but that is a design change and does not belong in a bugfix”。

- 请求字段 serialized_named_tensors 与 #7287 的分歧 (缺陷 3) (design): 按当前 slang 的字段构造 adapter 请求，与同文件基础权重同步保持同一形态，两个调用点对齐并随 slang 演进同步。
- wake_up() 对称性修复是否被证明必要 (question): 作者明确这是防御性修复而非跑批证明的必要修复；保留但如实说明，负向对照无法复现原错误。
- 单元测试缺失与端到端验证的取舍 (testing): 接受该策略；作者表示若 CI 能承载 SGLang LoRA e2e 任务可贡献。
- .base_layer 剥离是权宜修复而非根治 (design): 接受权宜修复，vLLM 路径不动；引擎按 rollout 后端分支的新设计明确延后到独立变更。
- 已知缺口：多桶适配器同步未验证 (testing): 未解决；需要更大显存与更大模型的环境补充验证。
- 代码格式请求 (style): 作者在提交清单中勾选已应用 pre-commit 检查，最终 wuxibin89 APPROVED。

风险与影响

- 风险：1) 线协议版本耦合：serialized_named_tensors 字段形态绑定 slang $\geq 0.5.18$ ，若环境锁定旧版 (#7289 报告的环境即 slang 0.5.8) 会反向不兼容，需要依赖同步。2) 无自动化测试：所有证据为手工 e2e，后续改动静默回归风险高，CI 无法拦截。3) 非 SGLang 后端误伤：engine_workers.py 用 getattr(self.rollout, 'sleep_level', 2) 读取 adapter 语义，依赖各后端 sleep_level 含义一致，其他后端若定义不同语义需复核（材料未覆盖，标注不确定性）。4) Megatron 路径证据薄弱：仅 0.6B/12 步 / 单一 TP 配置，未覆盖大模型与 TP 变化。5) 分桶同步未验证：update_weights_bucket_megabytes 多桶路径的 adapter 同步无任何证据（作者自述 Known gap）。6) 行为变更：正则 target_modules 从静默损坏变为启动报错，--lora-target-modules all 取代 all-linear，依赖旧行为的启动脚本需调整。7) .base_layer 字符串替换是 workaround，若真实权重名中恰含该段会被误删（概率低）。
- 影响：功能面是质变：LoRA + SGLang 组合在 FSDP 与 Megatron 两条训练路径首次端到端可用，补上了“LoRA 示例全用 vLLM”的盲区，用户不再被限制在 vLLM 后端做 LoRA 训练。系统面：改动集中在 SGLang rollout 与引擎同步路径，vLLM 路径与训练

核心算法不动，peft_config 新增的 peft_type 键对 vLLM 消费方是纯增量；非 LoRA 路径的唯一触碰点是 engine_workers 的 resume 条件，sleep_level 默认 2 时行为不变。团队面：为后续添加 SGLang LoRA 示例与 CI e2e 任务铺路，也暴露了多生产者共享一个 dict 契约缺少文档化的问题（已在 engine/base.py 补 docstring）。影响程度中高——范围收窄但对该组合是质变。

- 风险标记：无单测 /CI 覆盖，线协议依赖 sglang 版本，分桶同步未验证，Megatron 路径验证有限，关键路径改动

关联脉络

- PR #7287 [sglang] fix: make FSDP+LoRA adapter sync to SGLang work end-to-end: 本 PR 完整携带其缺陷 2 (asdict 误用) 的修复，并在缺陷 3 上修正其按 sglang 0.5.8 写的字段形态，与基础权重同步对齐。
- PR #7288 [sglang] fix: translate LoRA target_modules into SGLang's "all" sentinel: 缺陷 1 的修复，本 PR 原样携带；其 helper 与本 PR 新增的 sglang_lora_target_modules() 同源相邻。
- PR #7339 (原标题未提供) 针对缺陷 4 同一崩溃的守卫：PR body 提及 #7339 的守卫方向正确但只读 sleep_level，而该值在首次同步后才置 1，单独不够；本 PR 携带等价守卫并补齐构造期推导与 wake_up() 对称。
- PR #7289 LoRA adapter path resumes the weights tag that was never released: 缺陷 4 的根源 issue，本 PR 通过 engine_workers.py 跳过 resume 与构造期 sleep_level 推导将其关闭。
- PR #7290 The two engines return differently shaped peft_config from get_per_tensor_param: megatron peft_config 缺 peft_type 的 issue，本 PR 在 build_peft_config_for_vllm() 补上并关闭，同时在 engine/base.py 文档化契约。
- PR #7327 [vllm] fix: resolve .base_layer on the vLLM receiver for non-merged LoRA sync: 同一 .base_layer 权重名家族：本 PR 的 _strip_lora_base_layer() 是同一问题在 SGLang 接收侧的对称处理。
- PR #7376 [megatron] fix: per-name, mapper-aware .base_layer strip in resolve_weight_name: 同为 .base_layer 权重名解析问题，本 PR 的字符串替换是更朴素但同目的的做法。