

PR #32707 完整报告

sgl-project/sglang

Split #32584 into 1/2: [LoRA] Guard DP-attention idle forwards against stale LoRA batch state

合并时间: 2026-08-01 05:23

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/32707>

执行摘要

- 一句话: LoRA 空闲 DP 前向增加批次状态防护, 避免崩溃
- 推荐动作: 值得精读 `lora_active` 门控和 `reset_batch_state()` 的生命周期设计, 是一个“批次级元数据显式管理”的典型示例。建议结合 2/2 #32708 一起阅读完整方案, 并关注后续是否补充针对 idle forward 的单元测试。

功能与动机

PR body 明确描述: "Idle DP forwards skip `prepare_lora_batch()`, so LoRA kernels ran with the previous batch's metadata and crashed or corrupted output." 在 DP-attention 下, 部分 rank 本地 batch 为空 (idle forward), `ForwardBatch.init_new` 提前返回, LoRA 后端的 `batch_info` 残留上一批的段指针与长度, LoRA 内核按过期元数据对空输入运行, 造成越界崩溃或静默输出损坏。

实现拆解

1. 建立状态清理通路: `BaseLoRABackend.__init__` 显式把 `batch_info` 初始化为 `None`, 新增 `reset_batch_state()` 清空 `batch_info`、`lm_head_batch_info`、`lm_head_pass_batch_infos`、`_lm_head_pass_idx`; `TritonLoRABackend` 覆写该方法并同步清空 `sgemm_batch_info`; `LoRAManager` 新增 `reset_lora_batch()` 作为统一入口。
`batch_info is None` 成为“未准备批次”的主信号。
2. 挂接空闲分支: `ForwardBatch.init_new` 的 `is_idle()` 分支在 `server_args.enable_lora` 时调用 `model_runner.lora_manager.reset_lora_batch()`, 保证每个 idle 前向开局处于干净状态。
3. 层级别门控: `BaseLayerWithLoRA` 新增 `lora_active` property (`self.set_lora` and `self.lora_backend.batch_info is not None`) ; `VocabParallelEmbeddingWithLoRA`、`ParallelLMHeadWithLoRA`、`ColumnParallelLinearWithLoRA` 及其 merged/row 变体、MoE 包装层的 `forward` 全部从 `self.set_lora` 切换为 `self.lora_active`; `trtllm_lora_temp` 的 O7/O8/O9/O10/O11 旁路流前向同步切换, MoE 层在 `batch_info is None` 时直接委托 base 层。
4. 防御性契约: `TritonLoRABackend._sgemm_info()` 将由 `getattr` 的静默回退改为显式断言 `batch_info is not None`; `deepseek_mla_correction.py` 删除 `hasattr(lora_backend, "batch_info")` 检查。

5. 测试与配套：本 PR 未新增测试文件（PR body checklist 声称已加单测，但变更列表不含测试）；完整 dp-attention e2e 验证集中在 2/2 #32708。

关键文件：

- python/sglang/srt/lora/layers.py（模块 LoRA 层；类别 source；类型 core-logic；符号 lora_active）：核心门控所在：新增 lora_active property，并将各 LoRA 包装层 forward 从 set_lora 切换为 lora_active；同时透传 reduce_results 保证 DeepseekV2AttentionMLA 空闲断言成立。
- python/sglang/srt/lora/backend/base_backend.py（模块 LoRA 后端；类别 source；类型 core-logic；符号 reset_batch_state）：定义 batch_info=None 主信号与 reset_batch_state()，是所有 LoRA 后端状态清理的基类契约。
- python/sglang/srt/lora/backend/triton_backend.py（模块 LoRA 后端；类别 source；类型 core-logic；符号 reset_batch_state）：覆写 reset_batch_state() 清空 sgemm_batch_info，并把 _sgemm_info() 改为断言式防误用。
- python/sglang/srt/lora/lora_manager.py（模块 LoRA 管理；类别 source；类型 core-logic；符号 reset_lora_batch）：新增 reset_lora_batch() 统一入口，被 ForwardBatch 空闲分支调用。
- python/sglang/srt/model_executor/forward_batch_info.py（模块 前向构造；类别 source；类型 data-contract）：修复入口：idle 分支在启用 LoRA 时调用 reset_lora_batch()，是整条防护链的触发点。
- python/sglang/srt/lora/trtllm_lora_temp/attention.py（模块 LoRA 后端；类别 source；类型 core-logic）：O7/O8/O10/O11 旁路流前向的门控从 set_lora 切换为 lora_active，保证空闲前向退回原始 forward。
- python/sglang/srt/lora/deepseek_mla_correction.py（模块 MLA 修正；类别 source；类型 core-logic）：删除 hasattr(lora_backend, "batch_info") 防御，因 batch_info 已在构造时保证存在，简化状态读取。
- python/sglang/srt/lora/trtllm_lora_temp/merged_column.py（模块 LoRA 后端；类别 source；类型 core-logic）：O9 merged-column 旁路流前向同步切换为 lora_active 门控。

关键符号：lora_active, reset_batch_state, reset_lora_batch, _sgemm_info

关键源码片段

python/sglang/srt/lora/layers.py

核心门控所在：新增 lora_active property，并将各 LoRA 包装层 forward 从 set_lora 切换为 lora_active；同时透传 reduce_results 保证 DeepseekV2AttentionMLA 空闲断言成立。

文件：python/sglang/srt/lora/layers.py（节选）

```
class BaseLayerWithLoRA(nn.Module):
    def __init__(self, base_layer: nn.Module, lora_backend: BaseLoRABackend):
        super().__init__()
        self.base_layer = base_layer
        self.set_lora: bool = False
        self.lora_backend = lora_backend
```

```

if hasattr(self.base_layer, "weight"):
    self.weight = self.base_layer.weight
if hasattr(self.base_layer, "bias") and self.base_layer.bias is not None:
    self.bias = self.base_layer.bias
# 转发 base 层的 reduce_results: DP-attention 空闲前向下
# DeepseekV2AttentionMLA 会断言 o_proj.reduce_results 为假,
# LoRA 包装后该属性必须仍然可见
if hasattr(self.base_layer, "reduce_results"):
    self.reduce_results = self.base_layer.reduce_results

@property
def lora_active(self) -> bool:
    """LoRA 生效条件 = 层已配置 LoRA buffer 且当前前向有批次元数据。

    batch_info 在空闲前向中被 reset_batch_state() 清为 None,
    因此空闲前向自动走基础路径, 不会读到上一批的过期段指针。"""
    return self.set_lora and self.lora_backend.batch_info is not None

```

python/sglang/srt/lora/backend/base_backend.py

定义 `batch_info=None` 主信号与 `reset_batch_state()`, 是所有 LoRA 后端状态清理的基类契约。

文件: python/sglang/srt/lora/backend/base_backend.py (节选)

```

class BaseLoRABackend(LoRABackendLmHeadMixing):
    def __init__(self, max_loras_per_batch: int, device: torch.device):
        self.max_loras_per_batch = max_loras_per_batch
        self.device = device
        # batch_info 在 prepare_lora_batch() 中为每个前向设置;
        # 空闲前向由 reset_batch_state() 清空, None 表示“未准备批次”,
        # 层守卫 lora_active 依赖这个主信号决定是否应用 LoRA
        self.batch_info: Optional[LoRABatchInfo] = None
        self.init_lm_head_config()
        self._is_moe_lora = False
        self.prefill_cuda_graph_batch_info: LoRABatchInfo | None = None
        self.prefill_cuda_graph_max_bs: int | None = None
        self.prefill_cuda_graph_max_tokens: int | None = None

    def reset_batch_state(self):
        """prepare_lora_batch() 的空闲前向对应操作: 清空全部分批元数据。"""
        self.batch_info = None
        self.lm_head_batch_info = None
        self.lm_head_pass_batch_infos = None
        self._lm_head_pass_idx = None

```

python/sglang/srt/model_executor/forward_batch_info.py

修复入口: idle 分支在启用 LoRA 时调用 `reset_lora_batch()`, 是整条防护链的触发点。

文件: python/sglang/srt/model_executor/forward_batch_info.py (节选, ForwardBatch.init_new

内)

```
# idle 前向：本 rank 没有本地 token，直接构造空 positions 返回。
# DP-attention 下该分支会跳过 prepare_lora_batch()，因此必须显式
# 清空 LoRA 批次状态，否则 LoRA 内核会读上一批的过期元数据
if ret.forward_mode.is_idle():
    ret.positions = torch.empty((0,), dtype=torch.int64, device=device)
    if model_runner.server_args.enable_lora:
        model_runner.lora_manager.reset_lora_batch()
    return ret
```

评论区精华

来自 #32584 的 review 反馈（提交 2 信息回述）：重置 per-batch LoRA 状态应是独立操作，而不是 `prepare_lora_batch()` 的特殊情况。作者据此把 `idle` 重置改为独立方法 `reset_lora_batch()`，`prepare_lora_batch()` 不再感知空闲前向，职责更清晰。另一个亮点是 `_sgemm_info()` 的新增断言，把“未备批次就调 LoRA 内核”从静默行为变成显式失败，并提示调用方按 `lora_active` 门控。

- 空闲前向 LoRA 重置应独立成方法而非 `prepare_lora_batch` 特判 (design): 新增 `LoRAManager.reset_lora_batch()` → `BaseLoRABackend.reset_batch_state()`, `ForwardBatch` 空闲分支直接调用, `prepare_lora_batch()` 不再感知 `idle` 语义。

风险与影响

- 风险：门控切换面广：layers.py 与 trtllm_lora_temp 共 7 处 forward 从 `set_lora` 改为 `lora_active`，依赖 `batch_info` 生命周期严格配对；若预填充 CUDA 图、spec decode 等路径意外清空状态，可能静默跳过 LoRA 应用。`reset_batch_state()` 是与 `prepare_lora_batch()` 成对维护的契约，未来新增 per-batch 状态（尤其新后端）时容易漏清。本 PR 无直接单元测试，回归防护不足。非 DP-attention 路径 `batch_info` 在正常前向总会被设置，行为等价，风险集中在空闲分支与 CUDA 图捕获。
- 影响：影响面限于 `--enable-dp-attention + LoRA` 组合场景，修复后该场景不再崩溃或输出损坏；非 DP-attention LoRA 行为保持不变。性能上 `idle` 分支只多一次轻量方法调用，可忽略。对团队而言，本 PR 立下了“LoRA 批次状态必须成对清理”的契约，为 #32708 的动态 LoRA 端点和后续新后端接入提供了基础。
- 风险标记：核心路径变更，缺少测试覆盖，状态同步易遗漏，跨后端契约变更

关联脉络

- PR #32584 LoRA under `--enable-dp-attention` (原始 PR, 被拆分为 1/2 与 2/2)：本 PR 是 #32584 的拆分 1/2，动机与方案直接继承自它；提交 2 的 review 反馈也来自 #32584。
- PR #32708 Split #32584 into 2/2: LoRA DP-attention attn-TP sharding: PR body 指明 2/2 增加 attn-TP 分片并解锁动态 LoRA 端点，携带完整 dp-attention e2e 结果；本 PR 为其前置修复。