

PR #2320 完整报告

THUDM/slime

[cleanup] remove dead code and merge never visited branches

合并时间: 2026-08-24 21:16

原文链接: <http://prhub.com.cn/THUDM/slime/pull/2320>

执行摘要

- 一句话: 清理死代码与未访问分支, 精简 32 个文件
- 推荐动作: 该 PR 可作为仓库清理死代码的范式, 适合需要了解如何系统识别未访问分支的工程师阅读。不建议深入学习具体业务逻辑, 但可以重点关注其删除策略: 先搜索调用方, 再合并分支, 最后同步测试。同时提醒在删除公共 API 前确认是否存在插件或外部依赖。

功能与动机

PR 标题明确为 cleanup, 目的是删除死代码并合并从未被访问的分支。此类清理通常源于静态检查或人工 review 发现: 某些函数、参数或类路径在现有调用链中已无引用 (如 greedy_partition、chunk_named_params_by_size、_TensorBackuperNoop、HfWeightIteratorBase 等), 保留它们会增加维护负担并误导后续开发者。

实现拆解

1. 删除无引用的工具函数: 在 slime/utils/seqlen_balancing.py 中移除 greedy_partition 与 get_reverse_idx, 在 slime/utils/misc.py 中移除 chunk_named_params_by_size 与 _chunk_by_size, 在 slime/utils/http_utils.py 中移除 terminate_process, 在 slime/utils/health_monitor.py 中移除 is_checking_enabled。这些函数在当前代码库中均无调用方, 删除后不影响运行行为, 但减少了 copy、heapq 等历史依赖的引用。
2. 简化 TensorBackuper: slime/utils/tensor_backuper.py 从抽象工厂模式收敛为单一具体类。删除了 TensorBackuper.create 静态工厂、_TensorBackuperNoop 类以及 _compute_hash_dict / _compute_hash_tensor 辅助函数, 原有 _TensorBackuperNormal 的逻辑直接并入 TensorBackuper。这样所有调用方不再需要区分“正常备份”与“无需备份”两种模式, 心智负担显著降低。
3. 合并 named_params_and_buffers 分支: slime/backends/megatron_utils/update_weight/common.py 中将 named_params_and_buffers 的 convert_to_global_name 与 translate_gpu_to_cpu 参数移除, 只保留全局命名路径; 同时删除 _named_params_and_buffers_vanilla、_compute_fqn、_maybe_get_cpu_backup 等辅助实现。调用方 (如 hf_checkpoint_saver.py、hf_weight_iterator_direct.py) 同步简化, 不再传多余参数。
4. 精简 TrainProfiler 与相关配置: slime/observability/profile_utils.py 中删除 iterate_train_actor、iterate_train_log_probs 和 _profile_simple_loop, TrainProfiler 不再输出单独的 train_actor / train_log_probs 分析目标; 同时去掉 profile_target 对

train_overall 的过滤，只要开启 use_pytorch_profiler 就直接创建整体 profiler。
slime/utils/arguments.py 中对应的配置项或校验分支也同步收敛。

5. 删除冗余抽象基类与测试配套：slime/backends/megatron_utils/update_weight/hf_weight_iterator_base.py 整个文件被删除，HfWeightIteratorBase 的创建与迭代职责由 HfWeightIteratorDirect 直接承担。同时 slime/backends/megatron_utils/hf_checkpoint_saver.py 等调用方移除对基类的引用。仓库内多个测试文件同步更新，移除对已删除函数或参数的断言。

关键文件：

- slime/utils/tensor_backper.py（模块 张量备份；类别 source；类型 deletion；符号 TensorBackuper, _TensorBackuperNoop, create, backup_tags）：核心清理点：移除 abstractmethod 与 Noop 分支，TensorBackuper 从抽象工厂简化为单一具体类，直接暴露备份语义。
- slime/backends/megatron_utils/update_weight/hf_weight_iterator_base.py（模块 权重迭代；类别 source；类型 deletion；符号 HfWeightIteratorBase, create, get_hf_weight_chunks）：删除抽象基类 HfWeightIteratorBase，消除不必要的间接层，HfWeightIteratorDirect 成为唯一实现。
- slime/backends/megatron_utils/update_weight/common.py（模块 权重更新；类别 source；类型 core-logic；符号 named_params_and_buffers, _maybe_get_cpu_backup, _named_params_and_buffers_vanilla, _compute_fqn）：合并 named_params_and_buffers 的 never visited 分支：去掉 vanilla 命名路径与 GPU→CPU 备份转换，只保留 global 路径，简化调用契约。
- slime/observability/profile_utils.py（模块 性能剖析；类别 source；类型 core-logic；符号 TrainProfiler, iterate_train_actor, iterate_train_log_probs, _profile_simple_loop）：移除 train_actor/train_log_probs 独立 profiler 迭代器，并取消 profile_target 对 train_overall 的过滤，减少配置分支。
- slime/utils/seqlen_balancing.py（模块 序列均衡；类别 source；类型 deletion；符号 greedy_partition, get_reverse_idx）：删除未使用的 greedy_partition 与 get_reverse_idx，保留主入口 get_seqlen_balanced_partitions。
- slime/utils/misc.py（模块 工具函数；类别 source；类型 deletion；符号 chunk_named_params_by_size, _chunk_by_size）：删除未使用的 chunk_named_params_by_size 与 _chunk_by_size，清理通用工具函数。

关键符号：TensorBackuper.background, TensorBackuper.restore,
TensorBackuper.copy, named_params_and_buffers, TrainProfiler.init, TrainProfiler.step

关键源码片段

slime/utils/tensor_backper.py

核心清理点：移除 abstractmethod 与 Noop 分支，TensorBackuper 从抽象工厂简化为单一具体类，直接暴露备份语义。

```
from collections import defaultdict
from collections.abc import Callable, Iterable
```

```

import torch

from slime.utils import accelerator

# 定义 source_getter 的类型: 返回 ( 名称, Tensor) 的可迭代序列
_SourceGetter = Callable[[], Iterable[tuple[str, torch.Tensor]]]

class TensorBackuper:
    """张量备份器: 将参数/缓冲区快照到 CPU 固定内存, 支持多 tag 独立备份与恢复。
    清理前该类还有抽象工厂和 Noop 变体, 现统一为具体类, 仅保留必走路径。
    """

    def __init__(self, source_getter: _SourceGetter):
        # source_getter 用于获取当前参数或缓冲区的名称与张量
        self._source_getter = source_getter
        # 按 tag 组织备份字典, 默认值为空 dict
        self._backups: dict[str, dict[str, torch.Tensor]] = defaultdict(dict)

    @property
    def backup_tags(self):
        # 返回所有已备份的 tag 列表, 供外部查询
        return list(self._backups)

    def get(self, tag: str):
        # 获取指定 tag 下的备份字典 ({ 名称: CPU 张量 })
        return self._backups[tag]

    @torch.no_grad()
    def backup(self, tag: str) -> None:
        # 将当前所有参数异步拷贝到 CPU 固定内存, 非阻塞
        backup_dict = self._backups[tag]
        for name, param in self._source_getter():
            if name not in backup_dict:
                # 首次备份时分配 CPU 固定内存, 后续复用
                backup_dict[name] = torch.empty_like(param, device=torch.device('cpu'), pin_memory=True)
            backup_dict[name].copy_(param.detach(), non_blocking=True)
        accelerator.synchronize()

    @torch.no_grad()
    def copy(self, *, src_tag: str, dst_tag: str):
        # 将 src_tag 的备份整体拷贝到 dst_tag, 用于不同快照间的复制
        for name in self._backups[dst_tag]:
            self._backups[dst_tag][name].copy_(self._backups[src_tag][name])

    @torch.no_grad()
    def restore(self, tag: str) -> None:
        # 用指定 tag 的备份恢复当前参数, 需保证所有参数都有备份

```

```

backup_dict = self._backups[tag]
for name, param in self._source_getter():
    assert name in backup_dict
    param.copy_(backup_dict[name], non_blocking=True)
accelerator.synchronize()

```

slime/backends/megatron_utils/update_weight/common.py

合并 named_params_and_buffers 的 never visited 分支：去掉 vanilla 命名路径与 GPU→CPU 备份转换，只保留 global 路径，简化调用契约。

```

# 统一入口：生成全局一致命名的参数 / 缓冲区迭代器，跨 PP/EP 保持一致
def named_params_and_buffers(args: Namespace, model: Sequence[torch.nn.Module]) ->
Iterator[tuple[str, torch.Tensor]]:
    """
    Yield (global_name, param/buffer) with consistent names across PP/EP.
    调整虚拟 PP 与 EP 的索引偏移，处理 decoder.layers、mtp.layers 与 expert_bias。
    """

    ep_size = mpu.get_expert_model_parallel_world_size()
    ep_rank = mpu.get_expert_model_parallel_rank()
    if args.num_experts:
        # expert 的层偏移 = 当前 EP rank * num_experts / ep_size
        expert_offset = ep_rank * args.num_experts // ep_size

    sig = inspect.signature(get_transformer_layer_offset)
    need_vp_stage = 'vp_stage' in sig.parameters

    for vp_stage, model_module in enumerate(model):
        if need_vp_stage:
            layer_offset = get_transformer_layer_offset(model_module.config, vp_stage)
        else:
            layer_offset = get_transformer_layer_offset(model_module.config)
        for name, param in model_module.named_parameters():
            # 兼容未做 ddp wrap 的模型，补上 module. 前缀
            if not name.startswith('module.module.'):
                name = 'module.' + name
            # 然后根据 decoder.layers / mtp.layers 等 pattern 调整全局层号
            # 后续逻辑（省略正则细节）会把原始名称改写为全局名称后 yield

```

slime/observability/profile_utils.py

移除 train_actor/train_log_probs 独立 profiler 迭代器，并取消 profile_target 对 train_overall 的过滤，减少配置分支。

```

class TrainProfiler:
    def __init__(self, args):
        self.args = args
        self._torch_profiler_overall = None
        self._memory_profiler_overall = None

    # 只要开启 PyTorch profiler，就创建整体训练 profiler（不再要求 profile_target 包含 train_

```

```

overall)
if args.use_pytorch_profiler:
    self._torch_profiler_overall = _create_torch_profiler(args, name='train_overall')

# record_memory_history 开启时立即启动内存记录器
if args.record_memory_history:
    self._memory_profiler_overall = _BaseMemoryProfiler.create(args)
    self._memory_profiler_overall.start()

def on_init_end(self):
    # 初始化结束后启动整体 profiler，避免记录初始化阶段
    if self._torch_profiler_overall is not None:
        self._torch_profiler_overall.start()

def step(self, rollout_id: int):
    # 每个 rollout 步调用一次，推进 torch profiler 的 schedule
    if self._torch_profiler_overall is not None:
        self._torch_profiler_overall.step()

# 达到指定步数后停止内存快照记录，避免长时间占用内存
if (
    self._memory_profiler_overall is not None
    and ((s := self.args.memory_snapshot_num_steps) is not None)
    and (rollout_id == s - 1)
):
    self._memory_profiler_overall.stop()

```

评论区精华

该 PR 没有产生任何 review 评论或讨论线程，由作者直接合并。从变更内容看，作者主要依赖于静态检查和调用方搜索来判定死代码。

- 暂无高价值评论线程

风险与影响

- 风险：主要风险点集中在“死代码”判定是否准确，以及删除后是否遗漏外部调用方。
profile_utils.py 中去掉 profile_target 过滤后，只要 use_pytorch_profiler=True 就会创建 train_overall profiler，若存在仅当 profile_target 不含 train_overall 时运行的任务，可能引入额外 profiling 开销；common.py 移除 translate_gpu_to_cpu 后，依赖 torch_memory_saver 的 CPU 备份路径不再可用，需确认所有调用方均未使用该参数；tensor_backper.py 删除 Noop 分支后，若存在需要“只校验不备份”的场景，行为会发生改变。建议审计相关调用的历史代码，尤其是插件目录与外部脚本。
- 影响：本次变更涉及 slime/utils、slime/backends、slime/observability 等多个核心子包，但均为删除性变更，未引入新功能。对最终用户无感知，对内部开发者而言，API 契约收窄（例如 named_params_and_buffers 参数减少、TrainProfiler 方法删除），需要同步适配调用方。整体上提高了代码可维护性，减少了配置分支，属于正向净化。

- 风险标记：删除可能遗漏外部调用方，profiler 行为变化，公共 API 参数移除

关联脉络

- PR #2322 [cleanup] Remove rollout_validation.py: 同属 cleanup 系列，删除冗余模块，进一步清理 ray 路径。
- PR #2321 [cleanup] extract create_weight_updater to make actor's init func cleaner: 同属 cleanup 系列，精简 actor 初始化，与本次更新权重模块的简化方向一致。
- PR #2316 Remove megatron_patch for memory optimization: 删除不再需要的 megatron 补丁，与本次删除死代码的思路一致。
- PR #2298 [NFC] Add observability subfolder: 引入 observability 子模块，本次 profile_utils.py 的清理正是在该模块上的延续。