

PR #1924 完整报告

THUDM/slime

Reduce host memory with upgraded tms

合并时间: 2026-05-19 21:19

原文链接: <http://prhub.com.cn/THUDM/slime/pull/1924>

执行摘要

- 一句话: 升级 torch_memory_saver 减少主机内存占用
- 推荐动作: 推荐精读, 特别是 megatron.patch 中如何通过 TMS 的 region 替换现有缓冲区分配逻辑、以及 disable_grad_buffers_cpu_backup 与 nccl_ub 的互斥设计。该 PR 展示了深度集成第三方内存管理库以优化训练内存的典型模式。

功能与动机

减少训练时主机端 (CPU) 的内存占用, 通过升级 torch_memory_saver 并利用其零拷贝特性和禁用冗余梯度缓冲区备份来实现。PR body 引用自内部仓库的已有优化实践。

实现拆解

1. actor_group.py: 将 TMS 预加载动态库的搜索从单个固定文件名改为依次尝试两个文件名 (优先 CUDA 12 专用版本), 失败时抛出 FileNotFoundError, 提高跨 CUDA 版本的兼容性。
2. megatron.patch: 在 Megatron 的 _ParamAndGradBuffer 和 DistributedDataParallel 中新增 disable_grad_buffers_cpu_backup 参数, 当开启时使用 TMS 的 region 上下文管理器 (禁用 CPU 备份) 分配梯度缓冲区, 避免不必要的 CPU 内存占用; 同时添加与 nccl_ub 的互斥断言。
3. arguments.py: 在 slime_validate_args 中当 offload_train = True 时自动设置 disable_grad_buffers_cpu_backup = True; 并调用 reset_arg 重置 --record-memory-history 的默认值。
4. common.py: 在 _maybe_get_cpu_backup 中调用 torch_memory_saver.get_cpu_backup 时添加 zero_copy=True 参数, 避免额外 CPU 内存拷贝。
5. Dockerfile: 新增 TMS_CUDA_MAJOR 构建参数, 根据 CUDA 版本动态安装 TMS, 确保库文件命名正确。
6. version.txt: 更新镜像版本号。

关键文件:

- slime/ray/actor_group.py (模块 调度器; 类别 source; 类型 core-logic; 符号 _allocate_gpus_for_actor): 核心逻辑: 动态库搜索路径改进, 支持 CUDA 12 新命名的预加载库, 增强兼容性。

- docker/patch/latest/megatron.patch (模块 补丁; 类别 test; 类型 test-coverage; 符号 DistributedDataParallel, _ParamAndGradBuffer, _make_no_backup_context) :
Megatron 核心补丁: 引入 disable_grad_buffers_cpu_backup 选项, 允许在 TMS 环境下禁用梯度缓冲区 CPU 备份以节省主机内存。
- slime/utils/arguments.py (模块 参数; 类别 source; 类型 core-logic; 符号 slime_validate_args, add_debug_arguments) : 参数验证: 自动设置 disable_grad_buffers_cpu_backup, 并重置 --record-memory-history 的默认值。
- slime/backends/megatron_utils/update_weight/common.py (模块 后端; 类别 source; 类型 core-logic; 符号 _maybe_get_cpu_backup) : 零拷贝优化: 在获取 CPU 备份时使用 zero_copy=True, 避免额外内存拷贝。
- docker/Dockerfile (模块 部署脚本; 类别 infra; 类型 infrastructure) : 构建脚本: 新增 TMS_CUDA_MAJOR 参数以根据 CUDA 版本正确安装 torch_memory_saver。
- docker/version.txt (模块 部署脚本; 类别 docs; 类型 documentation) : 版本号更新至 nightly-dev-20260519a。

关键符号: _allocate_gpus_for_actor, _maybe_get_cpu_backup, slime_validate_args

评论区精华

该 PR 无 review 评论和讨论。

- 暂无高价值评论线程

风险与影响

- 风险:
 1. 依赖兼容性: 升级 TMS 可能引入与旧版本的接口不兼容, 但本 PR 已通过搜索多个文件名和添加 zero_copy 参数适配新版。
 2. Megatron 内部补丁: patch 修改了 Megatron 的 DDP 和缓冲区核心逻辑 (param_and_grad_buffer.py), 若未来 Megatron 升级可能导致补丁冲突。
 3. 配置默认值变更: 当 offload_train 启用时自动设置 disable_grad_buffers_cpu_backup, 可能影响未预期此行为的用户, 但该行为符合设计意图。
 4. 性能回退: 禁用 CPU 备份可能会在梯度同步失败时增加恢复难度, 但 TMS 的 region 机制已处理异常路径。- 影响: 用户影响: 所有使用 --offload_train 进行训练的用户将自动获得更低的主机内存占用 (约减少梯度缓冲区在 CPU 上的副本), 无需手动调整。系统影响: 需要安装特定版本的 torch_memory_saver (commit a193d9dd), Docker 构建会自动处理。团队影响: 后续维护需关注 Megatron 补丁与上游的同步。
- 风险标记: 依赖版本升级, Megatron 内部 patch, 新增配置参数

关联脉络

- PR #1916 [docker] update torch memory saver: 同样涉及 torch_memory_saver 的 Docker 升级, 本 PR 进一步扩展了其集成深度。

- PR #1882 fix ppo value offload bugs: 涉及 offload_train 场景，本 PR 优化了 offload 下的 CPU 内存使用。