

PR #2220 完整报告

THUDM/slime

Optimize update weight

合并时间: 2026-07-19 10:16

原文链接: <http://prhub.com.cn/THUDM/slime/pull/2220>

执行摘要

- 一句话: 优化 MoE 权重更新, 引入专家路由跳过非专家转换
- 推荐动作: 建议架构师和训练引擎维护者仔细阅读 `expert_routing.py` 中的拓扑差分检测和路由规划逻辑, 以及 `update_weight_from_tensor.py` 中如何组合选择性转换。测试部分仍可加强, 特别是多节点异构拓扑场景。

功能与动机

原有权重更新机制对所有参数统一进行 Megatron-to-HF 转换和广播, 在 MoE 模型上存在大量冗余通信和转换。通过引入专家路由, 可以识别每个专家参数应发送到的引擎 rank, 只传输必要的权重; 结合选择性转换进一步跳过非专家权重, 从而降低权重更新延迟和网络负载。

实现拆解

1. 新增专家路由模块(`slime/backends/megatron_utils/update_weight/expert_routing.py`): 定义 `_SGLangMoeTopology` 描述引擎并行拓扑, `_can_route_experts` 判断路由条件, `configure_expert_routing` 规划每个专家参数的目标 rank。
2. 集成到权重更新主流程(`update_weight_from_tensor.py`): 在 `__init__` 中获取 `_full_param_info_buckets`; `connect_rollout_engines` 增加 `engine_parallel_configs` 参数; 新增 `_build_flattened_tensor_data` 安全地构建平坦化张量; 新增 `_prepare_expert_weight_batch` 和 `_update_expert_weights` 处理专家权重发送。
3. 支持选择性转换(`hf_weight_iterator_direct.py`): `get_hf_weight_chunks` 新增 `should_convert_chunk` 参数, 允许跳过非专家 chunk 的 HF 转换; 提取 `pack_param_info_buckets` 复用。
4. 暴露引擎并行配置(`rollout.py`): `ServerGroup` 新增 `parallel_config()` 方法, 返回 `tp/pp/ep/moe_dp` 配置; `RolloutServer` 新增 `engine_parallel_configs` 属性透传到权重更新模块。
5. 接口适配(`hf_weight_iterator_base.py`, `hf_weight_iterator_bridge.py`): 抽象方法增加 `param_info_buckets` 参数, 桥接器抛出不支持异常。
6. 测试新增: `test_sglang_config.py` 验证并行配置推导, `test_empty_colocated_weight_bucket.py` 保障空权重桶场景。

关键文件:

- slime/backends/megatron_utils/update_weight/expert_routing.py (模块 权重更新; 类别 source; 类型 dependency-wiring; 符号 _ExpertParam, _ExpertTransfer, _SGLangMoeTopology, _config_value) : 新增模块, 实现专家路由核心逻辑, 包括拓扑检测、可行性判断和传输规划。
- slime/backends/megatron_utils/update_weight/update_weight_from_tensor.py (模块 权重更新; 类别 source; 类型 dependency-wiring; 符号 _build_flattened_tensor_data, _prepare_expert_weight_batch, _update_expert_weights) : 集成专家路由, 新增选择性转换和扁平化张量构建, 是权重更新主入口。
- slime/backends/megatron_utils/update_weight/hf_weight_iterator_direct.py (模块 权重更新; 类别 source; 类型 core-logic; 符号 get_hf_weight_chunks, _convert_to_hf_named_tensors, pack_param_info_buckets) : 核心迭代器, 支持选择性转换, 提取公共分桶函数。
- slime/ray/rollout.py (模块 调度器; 类别 source; 类型 core-logic; 符号 parallel_config, engine_parallel_configs) : 暴露引擎并行配置, 供权重更新模块使用。
- slime/backends/megatron_utils/update_weight/hf_weight_iterator_base.py (模块 权重更新; 类别 source; 类型 core-logic; 符号 get_hf_weight_chunks) : 抽象基类接口适配, 增加 param_info_buckets 参数。
- slime/backends/megatron_utils/update_weight/hf_weight_iterator_bridge.py (模块 权重更新; 类别 source; 类型 core-logic; 符号 get_hf_weight_chunks) : 桥接器适配接口, 抛出不支持 param_info_buckets 的异常。
- slime/backends/megatron_utils/update_weight/common.py (模块 权重更新; 类别 source; 类型 core-logic) : 公共函数调整, 支持新逻辑。
- slime/backends/sglang_utils/sglang_engine.py (模块 引擎; 类别 source; 类型 core-logic) : 引擎配置调整, 配合并行配置传递。
- slime/backends/megatron_utils/hf_checkpoint_saver.py (模块 检查点; 类别 source; 类型 core-logic) : 检查点保存适配配置变更。
- tests/utils/test_sglang_config.py (模块 测试; 类别 test; 类型 test-coverage; 符号 test_server_group_parallel_config_derives_tp_from_overridden_pp, test_sglang_server_args_derive_tp_from_overridden_pp) : 测试并行配置推导逻辑。
- tests/test_empty_colocated_weight_bucket.py (模块 测试; 类别 test; 类型 test-coverage) : 测试空 colocated 权重桶场景。
- slime/backends/megatron_utils/actor.py (模块 后端; 类别 source; 类型 core-logic) : 微小变更以兼容新配置。

关键符号: configure_expert_routing, _get_sglang_moe_topology, _get_homogeneous_sglang_moe_topology, _can_route_experts, _build_flattened_tensor_data, _prepare_expert_weight_batch, _update_expert_weights, get_hf_weight_chunks, _convert_to_hf_named_tensors, pack_param_info_buckets, parallel_config, engine_parallel_configs

关键源码片段

slime/backends/megatron_utils/update_weight/expert_routing.py

新增模块，实现专家路由核心逻辑，包括拓扑检测、可行性判断和传输规划。

注意： 以下代码定义 SGLang MoE 拓扑类型和相关检查函数

```
from dataclasses import dataclass
from typing import Any, Mapping, Sequence
from argparse import Namespace
```

```
@dataclass(frozen=True)
```

```
class _SGLangMoeTopology:
```

```
    tp_size: int
```

```
    pp_size: int
```

```
    ep_size: int
```

```
    moe_dp_size: int
```

```
def _config_value(
```

```
    parallel_config: Mapping[str, Any] | None,
```

```
    key: str,
```

```
    default: Any,
```

```
) -> Any:
```

```
    # 从 parallel_config 或 default 中取值
```

```
    if parallel_config is None:
```

```
        return default
```

```
    return parallel_config.get(key, parallel_config.get(key.replace('_', '-'), default))
```

```
def _get_sglang_moe_topology(
```

```
    args: Namespace,
```

```
    engine_gpu_count: int,
```

```
    parallel_config: Mapping[str, Any] | None = None,
```

```
) -> _SGLangMoeTopology:
```

```
    # 根据参数和引擎 GPU 数量计算拓扑
```

```
    pp_size = int(_config_value(parallel_config, 'pp_size', getattr(args, 'sglang_pp_size', 1)))
```

```
    default_tp_size = engine_gpu_count // pp_size
```

```
    tp_size = int(_config_value(parallel_config, 'tp_size', default_tp_size))
```

```
    ep_size = int(_config_value(parallel_config, 'ep_size', getattr(args, 'sglang_ep_size', 1)))
```

```
    moe_dp_size = int(_config_value(parallel_config, 'moe_dp_size', getattr(args, 'sglang_moe_dp_size', 1)))
```

```
    return _SGLangMoeTopology(tp_size, pp_size, ep_size, moe_dp_size)
```

```
def _can_route_experts(
```

```
    args: Namespace,
```

```
    sglang_moe_topology: _SGLangMoeTopology,
```

```
    engine_gpu_counts: Sequence[int],
```

```
) -> bool:
```

```
    # 检查是否可以进行专家路由：需要 raw 模式、PP=1、EP>1、无冗余专家等
```

```
    return (
```

```
        getattr(args, 'megatron_to_hf_mode', 'raw') == 'raw'
```

```
        and sglang_moe_topology.pp_size == 1
```

```

        and slang_moe_topology.ep_size > 1
        and not getattr(args, 'sglang_enable_eplb', False)
        and getattr(args, 'sglang_ep_num_redundant_experts', 0) == 0
        and getattr(args, 'sglang_init_expert_location', 'trivial') == 'trivial'
        and not getattr(args, 'sglang_enable_elastic_expert_backup', False)
        and mpu.get_expert_tensor_parallel_world_size() == 1
        and _sglang_moe_tp_is_one(engine_gpu_counts, slang_moe_topology)
    )

```

slime/backends/megatron_utils/update_weight/update_weight_from_tensor.py

集成专家路由，新增选择性转换和扁平化张量构建，是权重更新主入口。

注意：构建平坦化张量数据，避免重用 buffer 引发数据竞争

```

def _build_flattened_tensor_data(
    named_tensors: list[tuple[str, torch.Tensor]],
) -> dict[str, Any]:
    if not named_tensors:
        return {
            'flattened_tensor': torch.empty(0, dtype=torch.uint8, device=torch.cuda.current_device()),
            'metadata': [],
        }
    # 每次创建新的 FlattenedTensorBucket，确保不与前一次共享 buffer
    flattened_tensor_bucket = FlattenedTensorBucket(named_tensors=named_tensors)
    return {
        'flattened_tensor': flattened_tensor_bucket.get_flattened_tensor(),
        'metadata': flattened_tensor_bucket.get_metadata(),
    }

```

```

# 在 __init__ 中从 hf_weight_iterator 获取完整和专家参数桶
self._hf_weight_iterator = HfWeightIteratorBase.create(args, model, model_name, quantization_config)
param_info_buckets = getattr(self._hf_weight_iterator, 'megatron_local_param_info_buckets', None)
self._full_param_info_buckets = (
    tuple(tuple(bucket) for bucket in param_info_buckets) if param_info_buckets is not None else None
)
self._non_expert_param_info_buckets: list[list[ParamInfo]] | None = None

```

slime/ray/rollout.py

暴露引擎并行配置，供权重更新模块使用。

注意：根据 overrides 和 args 推导并行配置

```

def parallel_config(self) -> dict[str, Any]:
    overrides = {key.replace('-', '_'): value for key, value in self.sglang_overrides.items()}
    pp_size = int(overrides.get('pp_size', getattr(self.args, 'sglang_pp_size', 1)))

```

```

tp_size = int(overrides.get('tp_size', self.num_gpus_per_engine // pp_size))
return {
    'tp_size': tp_size,
    'pp_size': pp_size,
    'ep_size': int(overrides.get('ep_size', getattr(self.args, 'sglang_ep_size', 1))),
    'moe_dp_size': int(overrides.get('moe_dp_size', getattr(self.args, 'sglang_moe_dp_size', 1)
)),
}

```

```

@property
def engine_parallel_configs(self) -> list[dict[str, Any]]:
    """Per-engine SGLang parallel config, parallel to engines."""
    return [g.parallel_config() for g in self.server_groups for _ in g.engines]

```

评论区精华

- 暂无高价值评论线程

风险与影响

- 风险：专家路由的合法性依赖引擎拓扑的一致性假设，若 `engine_parallel_configs` 配置错误或引擎之间拓扑异构未被检测到，可能导致权重发送错误目标；`should_convert_chunk` 的判断条件若漏掉某些 `chunk` 会导致未转换的权重被发送；`_build_flattened_tensor_data` 每次创建新 `buffer` 增加了短暂内存峰值，但避免了并发写入风险。
- 影响：主要影响带有 MoE 的训练任务，权重更新性能提升；非 MoE 模型行为不变（路由条件不满足时回退旧路径）。需要调用方在 `connect_rollout_engines` 时额外传入 `engine_parallel_configs`，否则保留向后兼容。
- 风险标记：专家路由依赖拓扑一致，选择性转换可能遗漏参数，新增内存峰值

关联脉络

- PR #2089 [2/n] Disaggregated rollout: disk-level delta weight sync: 均为权重更新优化系列，2089 引入磁盘级 delta 同步，本 PR 进一步精细化专家路由。
- PR #2185 Support routed_experts_start_len: 涉及 MoE 专家路由相关的参数支持，为本 PR 的专家路由功能提供前置条件。
- PR #2181 [3/n] Disaggregated rollout: engine-side /pull_weights: 权重同步链路上的另一优化，本 PR 的专家路由可与 /pull_weights 协同。