

PR #7374 完整报告

verl-project/verl

[ci, trainer] feat: remove mindspeedllm backend engine support.

合并时间: 2026-08-14 14:04

原文链接: <http://prhub.com.cn/verl-project/verl/pull/7374>

执行摘要

- 一句话: 移除 Mindspeed-LLM 后端引擎 (约 1941 行)
- 推荐动作: 值得快速浏览 (10 分钟内可读完): 这是理解 verl 引擎注册架构的很好样例。重点关注 transformer_impl.py 中保留类与删除类的注册差异 (backend='megatron' vs backend='mindspeed_megatron'), 以及 constants_ppo.py 中运行时环境变量的收敛方式。若团队维护 NPU 训练配置, 务必确认没有遗留 mindspeed 策略引用, 并规划到 megatron 后端的迁移。整体改动机械、无评审争议, 不需要精读源码细节。

功能与动机

PRbody只有一句话: "removemindspeedllmbackendenginesupport.", 未关联任何issue。从代码结构可以还原动机: 仓库长期存在两条 NPU 引擎路径——其一是保留的 MindspeedEngineWithLMHead / MindspeedEngineWithValueHead (在 EngineRegistry 中以 backend="megatron"、device="npu" 注册, 通过 mindspeed.megatron_adaptor.repatch 为 Megatron 打补丁); 其二是本次删除的 MindSpeedMegatronEngineWithLMHead (backend="mindspeed_megatron", 依赖 mindspeed_llm.tasks.megatron_adaptor_v2.repatch 与独立的 gpt_model_provider)。后者是早期 NPU 支持引入的独立实现, 与新架构功能高度重叠、维护成本高; 而近期 PR#7345、PR#7293、PR#7337 等 NPU CI 工作正在持续收敛 Ascend 支撑面, 删除该后端正是这一收敛方向的自然延续。

实现拆解

1. 删除独立后端引擎实现

- verl/workers/engine/mindspeed/transformer_impl.py: 删除 MindSpeedMegatronEngineWithLMHead 类 (约 65 行)。该类在 _init_device_mesh 中调用 apply_patch, 在 _build_megatron_module 中通过 gpt_model_provider 构建模型, 并强制校验 vanilla_bridge=True; 保留的 MindspeedEngineWithLMHead / MindspeedEngineWithValueHead 不再需要这些接线, 只保留 _mindspeed_repatch 一个补丁入口。
- verl/workers/engine/mindspeed/utils.py: 删除 210 行, 包括 apply_patch、gpt_model_provider、set_global_config、add_mcore_arguments 与三个 get_base_mcore_config_from_* 转换函数; 仅保留 reset_fp8_reuse_quantized_weight, 供保留引擎的 to() 方法在设备迁移前处理 FP8 量化权重缓存。

2. 删除配置契约

- verl/workers/config/engine.py: 移除 MindSpeedEngineConfig (含 strategy 取值校验、dtype 校验、TP=1 时强制关闭 sequence_parallel 的逻辑), 并从 __all__ 导出中清除。
- verl/workers/config/actor.py 与 critic.py: 删除 MindSpeedActorConfig、MindSpeedCriticConfig 及其挂接 engine / checkpoint 的字段。
- verl/workers/config/checkpoint.py: 删除 MindSpeedCheckpointConfig (原为 McoreCheckpointConfig 的子类, 用于保留 mbridge 配置位)。

3. 收敛 trainer 侧运行时环境逻辑

- verl/trainer/constants_ppo.py: 删除 _uses_mindspeed 辅助函数, 以及 get_ppo_ray_runtime_env 中为 mindspeed 策略设置 CUDA_DEVICE_MAX_CONNECTIONS=1 的分支; 此后该环境变量只服务于 Megatron (Hopper / Ampere) 路径。

4. 同步删除配置、示例与测试

- 删除 4 个 Hydra 配置: verl/trainer/config/engine/mindspeed.yaml、actor/mindspeed_actor.yaml、ref/mindspeed_ref.yaml、critic/mindspeed_critic.yaml。
- 删除 2 个示例脚本 (examples/ascend_extras/grpo_trainer/ 下的 30B-A3B 与 32B MindSpeed) 和 3 个 NPU 测试脚本 (tests/special_npu/ 下的 8B、30B GRPO 及 nightly_ci_ascend 8B), nightly 脚本的删除直接缩小了 nightly_ascend.yml 的 job 面。

5. 更新 CI 与导入接线

- .github/workflows/e2e_ascend.yml、nightly_ascend.yml 分别删除约 70、68 行 mindspeedllm 相关 job。
- verl/workers/engine/init.py 与 verl/workers/engine/mindspeed/init.py 调整导入, 确保 EngineRegistry 不再注册 mindspeed_megatron 后端; 第二个 commit 为 Merge branch 'main' into llm, 说明合入前已与主干对齐。

上述五步围绕同一目标: 让重复的 mindspeed_megatron 后端从注册表、配置面、CI 面同时消失, 避免残留引用。

关键文件:

- verl/workers/engine/mindspeed/transformer_impl.py (模块引擎实现; 类别 source; 类型 core-logic; 符号 MindSpeedMegatronEngineWithLMHead, _mindspeed_repatch, _init_device_mesh, _build_megatron_module): 移除独立注册的 mindspeed_megatron 后端引擎类 MindSpeedMegatronEngineWithLMHead (含 apply_patch / gpt_model_provider 接线), 保留 backend='megatron' 的 MindSpeed repatch 路径, 是本次变更的核心。
- verl/workers/engine/mindspeed/utils.py (模块引擎实现; 类别 source; 类型 dependency-wiring; 符号 apply_patch, gpt_model_provider, set_global_config, add_mcore_arguments): 210 行配置转换与 Megatron 补丁工具函数被删除, 仅保留 FP8 缓存处理函数, 是删除面最大的源码文件。

- verl/workers/config/engine.py (模块 配置层; 类别 source; 类型 core-logic; 符号 MindSpeedEngineConfig, post_init) : 删除 MindSpeedEngineConfig 及策略 / 数据类型校验逻辑, 属于配置契约层面的破坏性变更。
- verl/workers/config/actor.py (模块 配置层; 类别 source; 类型 core-logic; 符号 MindSpeedActorConfig, post_init) : 删除 MindSpeedActorConfig, 清理 actor 侧 mindspeed 引擎与 checkpoint 接线。
- verl/trainer/constants_ppo.py (模块 训练器; 类别 source; 类型 core-logic; 符号 _uses_mindspeed, get_ppo_ray_runtime_env) : 删除 _uses_mindspeed 与 Ray 运行时环境中对应的 CUDA_DEVICE_MAX_CONNECTIONS 分支, 收敛环境变量逻辑。
- verl/trainer/config/engine/mindspeed.yaml (模块 配置层; 类别 config; 类型 deletion) : mindspeed 引擎 Hydra 配置删除, 与 dataclass 移除配套, 是用户侧配置破坏点。
- tests/special_npu/run_qwen3_30b_grpo_mindspeedllm.sh (模块 测试脚本; 类别 test; 类型 deletion) : 30B GRPO 端到端脚本被删除, 代表该后端的回归测试面消失。
- .github/workflows/e2e_ascend.yml (模块 工作流; 类别 infra; 类型 infrastructure) : e2e_ascend.yml 中约 70 行 mindspeedllm job 被移除, CI 面收敛。

关键符号: apply_patch, gpt_model_provider, set_global_config, add_mcore_arguments, get_base_mcore_config_from_model_config, get_base_mcore_config_from_engine_config, get_base_mcore_config_from_optim_config, MindSpeedMegatronEngineWithLMHead, _uses_mindspeed, MindSpeedEngineConfig, MindSpeedActorConfig, MindSpeedCriticConfig, MindSpeedCheckpointConfig

关键源码片段

verl/workers/engine/mindspeed/transformer_impl.py

移除独立注册的 mindspeed_megatron 后端引擎类 MindSpeedMegatronEngineWithLMHead (含 apply_patch / gpt_model_provider 接线), 保留 backend='megatron' 的 MindSpeed repatch 路径, 是本次变更的核心。

```
# verl/workers/engine/mindspeed/transformer_impl.py
# 本 PR 删除独立后端 MindSpeedMegatronEngineWithLMHead (backend = "mindspeed_megatron") 后,
# 保留文件的核心逻辑如下: NPU 上统一走 "megatron" backend + MindSpeed repatch 补丁。
```

```
try:
    from mindspeed.megatron_adaptor import repatch
except ImportError:
    repatch = None
```

```
def _mindspeed_repatch(engine_config):
    """对 Megatron 应用 MindSpeed 的运行补丁。
```

必须在 initialize_model_parallel 之前执行, 否则 CP 的 ring-rank 初始化包装器不会被注册: verl 通过 hydra 配置传入 CP size 而非 CLI 参数, 第一轮补丁时 context_parallel_size 仍是默认值 1。

```

"""
if repatch is not None:
    from verl.utils.megatron_utils import mapping_string_to_attn_backend

    repatch_config = mapping_string_to_attn_backend(dict(engine_config.get("override_
transformer_config", {})))
    # flash-attn-npu batch-invariant 会替换 DotProductAttention.forward;
    # fusion attention 在 use_flash_attn = True 时注册同一补丁,
    # 造成 "the patch of forward exist" 冲突, 因此二者需要互斥。
    if repatch_config.get("use_flash_attn_npu_batch_invariant"):
        repatch_config["use_flash_attn"] = False
    else:
        repatch_config.setdefault("use_flash_attn", True)
    if engine_config.context_parallel_size > 1:
        repatch_config["context_parallel_size"] = engine_config.context_parallel_size
    repatch(repatch_config)

@EngineRegistry.register(model_type="language_model", backend="megatron", device="npu")
class MindspeedEngineWithLMHead(MegatronEngineWithLMHead):
    # 保留的 NPU language model 引擎: 不再存在独立的 mindspeed_megatron 后端,
    # 全部复用 Megatron 引擎实现, 仅在设备网格初始化前插入 MindSpeed 补丁。

    def to(self, device: str, model: bool = True, optimizer: bool = True, grad: bool = True):
        """迁移参数 / 优化器状态前, 先处理 MindSpeed 的 FP8 量化缓存。

        若不先清理, NPU 上 fp8_reuse_quantized_weight 的旧缓存会与
        设备迁移后的权重不一致; release_storage = True 会释放量化权重缓存。
        """
        reset_fp8_reuse_quantized_weight(self, device, model, optimizer, grad)
        super().to(device=device, model=model, optimizer=optimizer, grad=grad)

@EngineRegistry.register(model_type="value_model", backend="megatron", device="npu")
class MindspeedEngineWithValueHead(MegatronEngineWithValueHead):
    def _init_device_mesh(self):
        # repatch 必须在 initialize_model_parallel 之前完成,
        # 保证 CP ring-rank 初始化包装器在第一次调用时就生效。
        _mindspeed_repatch(self.engine_config)
        super()._init_device_mesh()

```

verl/workers/engine/mindspeed/utils.py

210 行配置转换与 Megatron 补丁工具函数被删除, 仅保留 FP8 缓存处理函数, 是删除面最大的源码文件。

```

# verl/workers/engine/mindspeed/utils.py (本 PR 删除 210 行后的全部内容)
# 被删除的 apply_patch、gpt_model_provider、set_global_config、add_mcore_arguments
# 以及三个 get_base_mcore_config_from_* 转换函数, 均只服务于 legacy 的
# mindspeed_megatron 后端; 保留的该函数服务于 backend = "megatron" 的 NPU 引擎。

```

```
def reset_fp8_reuse_quantized_weight(engine, device: str, model: bool, optimizer: bool, grad:
bool):
    override_config = getattr(engine.engine_config, "override_transformer_config", None)
    # 仅在开启 FP8 量化权重重复用时才需要处理缓存
    if override_config and override_config.get("fp8_reuse_quantized_weight", False):
        from mindspeed.te.pytorch.fp8.reuse import (
            clear_weight_quantization_reuse_cache,
            set_weight_release_enabled,
        )

        # 在 NPU 上清理量化权重缓存，避免跨设备迁移后复用旧量化结果
        clear_weight_quantization_reuse_cache(release_storage=True)

        # 只有训练模式下的模块允许释放高精度权重：ref 模型需要保留高精度权重
        # 用于 offload，actor_update 模型则可在释放后于 optimizer step 前恢复
        set_weight_release_enabled(getattr(engine, "mode", None) == "train")
```

评论区精华

该 PR 的讨论记录为零：comments_count=0、review_comments_count=0、无关联 issue，wuxibin89 仅给出 APPROVED，未附带说明。这种 "零讨论直接合入" 在纯删除型清理中常见——保留路径与删除路径功能重叠度高，审阅者判断回归风险较低。但正因没有评审交流，PR body 缺失的动机细节（例如为何选在这个时点清理、是否与 mindspeed_llm 上游维护状况相关）只能从同仓库近期 NPU CI 收敛 PR 中推断，建议必要时向作者 pengnuoheng 求证。

- 暂无高价值评论线程

风险与影响

- 风险：配置破坏性变更：strategy='mindspeed' / 'mindspeed_megatron' 的 dataclass 与 yaml 全部消失，存量用户配置会直接报错；且本 PR 未同步更新 docs/ascend_tutorial 等文档，文档或示例中可能残留 mindspeed_megatron 引用（升级前建议 grep 排查）。测试覆盖缺口：3 个端到端脚本（含 nightly 8B）被删且无等价替代，保留的 NPU Megatron 路径（MindSpeedEngineWithLMHead）的端到端回归保障减弱，后续 NPU 改动将更依赖 e2e_ascend.yml 中的剩余 job。依赖残留：transformer_impl.py 与引擎 __init__ 仍以 try/except 方式导入 mindspeed，utils.py 仅剩的 reset_fp8_reuse_quantized_weight 仍强依赖 mindspeed.te.fp8.reuse，说明保留路径依然以 NPU 为前提，非 NPU 环境加载该模块的行为没有变化。迁移等价性未说明：原 30B-A3B 脚本依赖 mindspeed_llm.tasks.models.spec.qwen3_spec 等 mcore_kwargs（MoE 相关配置），迁移到 megatron 后端后这些参数如何映射未被本次变更说明，MoE 用户需要自行验证配置等价性。
- 影响：用户影响：仅影响 NPU 上显式使用 mindspeed / mindspeed_megatron 策略的用户，需迁移到 megatron 后端；Megatron / FSDP 等默认路径完全无感。系统影响：删除约 1941 行重复实现，EngineRegistry 注册表减少一个后端，NPU 训练链路收敛为单一 "megatron backend + MindSpeed repatch" 路径，后续 NPU 相关的 Megatron 修复（如 PR#7372 的 TND mask 修复）只需维护一条链路。团队影响：与 PR#7345、PR#7293、

PR#7337 等 NPU CI 收敛动作形成连贯演进，Ascend 支撑的维护成本显著下降；但该PR 无配套文档说明迁移方法，团队内部需要口头或文档补齐。

- 风险标记：配置破坏性变更，测试覆盖移除，遗留引用风险，NPU 路径迁移

关联脉络

- PR #7334 [misc] fix: warn on engine backend import failure instead of silent skip: 与本 PR 同改 verl/workers/engine/init.py，说明引擎注册 / 导入链路近期被持续打磨；本 PR 进一步让该文件不再介入 mindspeedllm 后端。
- PR #7345 [ci] chore: Fix npu nightly ci: 同时改动 nightly_ascend.yml 与 tests/special_npu/nightly_ci_ascend，本 PR 删除其中 mindspeedllm 的 8B nightly 脚本，属于同一 NPU CI 收敛路径。
- PR #7293 [ci] chore: Update ci image: 维护 Ascend CI 镜像与 nightly 调度，与本次 workflow 中移除 mindspeedllm job 同属 NPU CI 基础设施演进。
- PR #7337 [ci] chore: add three baselines for npu's nightly ci: 同为 nightly_ascend.yml 的演进，所加基线（REINFORCE++ 多轮等）部分替代了被删 mindspeedllm 脚本的覆盖职责。
- PR #7404 [ci] chore: Add npu ci env : NPU CI 环境变量与安装脚本的持续维护，与本次删除 mindspeedllm 遗留 job 呼应，说明 NPU 支撑正在聚焦单一路径。