

PR #45924 完整报告

vllm-project/vllm

[MoE Backend] add HPC-Ops MoE backend

合并时间: 2026-06-27 11:18

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/45924>

执行摘要

- 一句话: 集成 HPC-Ops FP8 MoE 后端, 适配 Hopper GPU
- 推荐动作: 值得精读, 特别是模块化 MoE 后端的注册机制和兼容层设计 (`_supports_*` 模式)。对于使用 H20/H200 的用户, 建议在 FP8 模型上评估性能后考虑启用。注意该后端依赖外部库, 部署时需确保环境安装正确。

功能与动机

HPC-Ops 是腾讯混元团队开发的高性能算子库, 在 H20 GPU 上对 FP8 MoE 有显著性能优势 (见 PR 中 benchmark)。vLLM 通过集成该后端, 为使用 H20 等 Hopper GPU 的中国用户提供更优的推理性能。PR body 指出 'Support hpc-ops moe backend', youkaichao 评论 'this is mainly for H20 GPUs used in china.'

实现拆解

1. 添加兼容层: 新增 `vllm/utils/hpc.py`, 通过 `has_hpc()` 检测 hpc 包是否安装, 并封装 `fuse_moe_impl` 和 `fuse_moe_blockwise_impl` 函数, 隔离直接依赖。
2. 实现模块化后端类: 新增 `vllm/model_executor/layers/fused_moe/hpc_moe.py`, 定义 `HPCExperts` 继承 `FusedMoEExpertsModular`, 实现设备检查、量化方案支持 (per-tensor 和 block-wise)、激活函数 (SiLU)、形状约束 (`hidden_dim % 128 == 0`) 及 workspace 计算。
3. 注册后端: 在 `vllm/model_executor/layers/fused_moe/oracle/fp8.py` 添加 `Fp8MoeBackend.HPC` 枚举, 在 `_get_priority_backends`、`backend_to_kernel_cls`、`map_fp8_backend` 和 `convert_to_fp8_moe_kernel_format` 中注册。
4. 配置支持: 在 `vllm/config/kernel.py` 的 `MoEBackend Literal` 中添加 "hpc", 并更新 docstring。
5. 文档记录: 在 `docs/design/moe_kernel_features.md` 表中添加 hpc 后端行。

根据 reviewer 建议, 移除了 benchmark 和测试文件, 仅保留核心代码和文档。

关键文件:

- `vllm/model_executor/layers/fused_moe/hpc_moe.py` (模块 MoE 执行器; 类别 source; 类型 data-contract; 符号 `HPCExperts`, `init`, `expects_unquantized_inputs`, `activation_format`): 核心实现, `HPCExperts` 模块化后端类, 声明支持条件和计算逻辑。

- `vllm/utils/hpc.py` (模块 工具层; 类别 `source`; 类型 `core-logic`; 符号 `has_hpc`, `fuse_moe_impl`, `fuse_moe_impl_fake`, `hpc_fuse_moe`) : HPC 包兼容层, 提供依赖检测和融合 MoE 函数, 隔离外部库。
- `vllm/model_executor/layers/fused_moe/oracle/fp8.py` (模块 后端选择器; 类别 `source`; 类型 `data-contract`) : 注册 HPC 后端到枚举、优先级列表和 kernel 映射。
- `vllm/config/kernel.py` (模块 配置层; 类别 `source`; 类型 `core-logic`) : 配置支持, 在 `MoEBackend` 字面量中添加 `"hpc"`。
- `docs/design/moe_kernel_features.md` (模块 文档; 类别 `docs`; 类型 `documentation`) : 文档记录新后端特性。

关键符号: `has_hpc`, `fuse_moe_impl`, `hpc_fuse_moe`, `fuse_moe_blockwise_impl`, `hpc_fuse_moe_blockwise`, `HPCExperts.init`, `HPCExperts._supports_current_device`, `HPCExperts._supports_quant_scheme`, `HPCExperts._supports_shape`, `HPCExperts.workspace_shapes`, `map_fp8_backend`, `backend_to_kernel_cls`

关键源码片段

`vllm/model_executor/layers/fused_moe/hpc_moe.py`

核心实现, `HPCExperts` 模块化后端类, 声明支持条件和计算逻辑。

```
# vllm/model_executor/layers/fused_moe/hpc_moe.py

class HPCExperts(mk.FusedMoEEpertsModular):
    """HPC-Ops 后端的 MoE 专家实现, 仅支持 Hopper FP8."""

    def __init__(
        self,
        moe_config: mk.FusedMoEConfig,
        quant_config: FusedMoEQuantConfig,
    ):
        super().__init__(moe_config, quant_config)
        # 断言仅 fp8 量化
        assert quant_config.weight_quant_dtype in (torch.float8_e4m3fn,), (
            "Only fp8 quantization is currently supported."
        )
        self.device = moe_config.device
        self.num_experts = moe_config.num_local_experts
        self.ep_rank = moe_config.moe_parallel_config.ep_rank
        self.ep_size = moe_config.moe_parallel_config.ep_size
        self.tp_rank = moe_config.moe_parallel_config.tp_rank
        self.tp_size = moe_config.moe_parallel_config.tp_size
        self.out_dtype = moe_config.in_dtype

    @staticmethod
    def _supports_current_device() -> bool:
        # 仅当 CUDA 且计算能力 >= 9.0 且 hpc 包可用时返回 True
        p = current_platform
```

```

    return (
        p.is_cuda()
        and (p.is_device_capability(90) or p.is_device_capability_family(100))
        and has_hpc()
    )

@staticmethod
def _supports_quant_scheme(
    weight_key: QuantKey | None,
    activation_key: QuantKey | None,
) -> bool:
    # 支持的量化方案: fp8 per-tensor 或 per-128 block-wise
    scheme = (weight_key, activation_key)
    return scheme in [
        (kFp8StaticTensorSym, kFp8StaticTensorSym), # per-tensor
        (kFp8Static128BlockSym, kFp8Dynamic128Sym), # block-wise G128
    ]

@staticmethod
def _supports_shape(hidden_dim: int) -> bool:
    # HPC 内核以 128 为 block 处理 hidden_size, 要求对齐
    return hidden_dim % 128 == 0

```

vllm/utils/hpc.py

HPC 包兼容层, 提供依赖检测和熔合 MoE 函数, 隔离外部库。

```
# vllm/utils/hpc.py
```

```

import functools
import importlib
import importlib.util
import torch

```

```
logger = init_logger(__name__)
```

```

@functools.cache
def has_hpc() -> bool:
    """检查 hpc 包是否可用, 避免潜在 CUDA 初始化。"""
    if importlib.util.find_spec("hpc") is None:
        logger.warning_once(
            "HPC attention requires the hpc module to be installed. "
            "Please install it from https://github.com/Tencent/hpc-ops"
        )
        return False
    return True

```

```

def fuse_moe_impl(
    x: torch.Tensor,
    gate_up_weight: torch.Tensor,

```

```

down_weight: torch.Tensor,
gate_up_scale: torch.Tensor,
down_scale: torch.Tensor,
act_and_mul_scale: torch.Tensor,
topk_ids: torch.Tensor,
topk_scale: torch.Tensor,
rank_ep: int,
num_expert_total: int,
use_bf16_mul: bool = True,
shared_output: torch.Tensor = None,
output: torch.Tensor = None,
) -> torch.Tensor:
    # 延迟导入 hpc 包, 仅在使用时加载
    from hpc import fuse_moe as fuse_moe_
    return fuse_moe_(
        x, gate_up_weight, down_weight,
        gate_up_scale, down_scale, act_and_mul_scale,
        topk_ids, topk_scale,
        rank_ep, num_expert_total,
        use_bf16_mul, shared_output, output=output,
    )

def hpc_fuse_moe(
    x: torch.Tensor,
    gate_up_weight: torch.Tensor,
    down_weight: torch.Tensor,
    gate_up_scale: torch.Tensor,
    down_scale: torch.Tensor,
    act_and_mul_scale: torch.Tensor,
    topk_ids: torch.Tensor,
    topk_scale: torch.Tensor,
    rank_ep: int,
    num_expert_total: int,
    use_bf16_mul: bool = True,
    shared_output: torch.Tensor = None,
    output: torch.Tensor = None,
) -> torch.Tensor:
    return fuse_moe_impl(
        x, gate_up_weight, down_weight,
        gate_up_scale, down_scale, act_and_mul_scale,
        topk_ids, topk_scale,
        rank_ep, num_expert_total,
        use_bf16_mul, shared_output, output=output,
    )

```

评论区精华

- robertgshaw2-redhat 询问形状约束和是否应设为默认后端, youkaichao 回应主要用于 H20 GPU, 不设默认。

- youkaichao 要求移除测试和 benchmark 代码，由 hpc-ops 团队内部测试，作者遵从。
- 多个 review 简化了代码：移除冗余设备能力检查和不必要注释，增加后端使用说明。
- yiaikwy-xpu-ml-framework-team 询问 HPC-Ops 在 FP32 上性能优于其他后端的具体原因，尚未回复，该疑问未解决。
- 形状约束与后端默认化 (question): youkaichao 回复 'this is mainly for H20 GPUs used in china.' 未设为默认。
- 移除冗余设备能力检查 (style): thisjiang 回复 'done' 并移除。
- 移除 Oracle 中多余注释 (style): thisjiang 删除注释。
- 增加后端适用场景描述 (documentation): thisjiang 在类文档字符串中添加说明。
- 移除 benchmark 和 test 文件 (testing): 作者移除相关文件，commit 记录显示删除 benchmark 和 test 代码。
- HPC-Ops 性能优于其他后端的原因 (question): 未收到回复，疑问未解决。

风险与影响

- 风险：
 1. 第三方依赖风险：运行时需要安装 hpc 包 (pip install hpc-ops)，若未安装，强制使用 `--moe_backend hpc` 可能导致 ImportError；尽管 `_supports_current_device` 会检查 `has_hpc()`，但用户强制指定时可能绕过该检查。
 2. 硬件限制：仅支持 SM90+ (Hopper)，在非 Hopper GPU 上使用可能失败（取决于 hpc 包的兼容性）。
 3. 量化方案限制：仅支持 fp8 per-tensor 和 block-wise (G128)，其他量化方案下即使指定 hpc 后端也可能运行出错。
 4. 缺少单元测试：测试代码已被移除，回归风险由外部团队承担，vLLM 侧无直接防护。 - 影响：对用户：新增 `--moe_backend hpc` 选项，默认 auto 不影响现有用户；H20 用户可手动启用获得更好性能。对系统：增加约 450 行代码，但通过模块化框架隔离，维护成本可控。对团队：需要与 hpc-ops 团队协作维护兼容性。 - 风险标记：第三方依赖，硬件限制，量化约束，缺少单元测试

关联脉络

- PR #46661 Allow FlashInfer A2A backends for TRTLLM FP8 MoE Modular: 同为新增 FP8 MoE 后端，注册到相同的 modular 框架。
- PR #45182 [Perf] Integrate TRTLLM BF16 MoE Modular Kernel: 类似的 MoE 后端集成 PR，可参考其注册方式。