

PR #44201 完整报告

vllm-project/vllm

[CPU][Zen] Route BF16 MoE inference through zentorch on AMD

合并时间: 2026-08-11 14:20

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/44201>

执行摘要

- 一句话: AMD Zen CPU 的 BF16 MoE 改走 zentorch 融合内核并带优雅回退
- 推荐动作: 值得精读。核心看点: (1) 两级能力检测设计——config 级 (`is_zentorch_moe_config_supported`) 与 layer 级 (`is_zentorch_moe_supported`) 分别服务 `is_supported_config` 与权重加载阶段, 职责清晰; (2) “跳过 prepack”所隐含的权重布局契约, 是接入外部内核时最容易踩坑的点; (3) 从旧 CPUFusedMOE 移植到 CPUUnquantizedExperts 的 rebase 过程, 展示了在模块化内核架构下如何平滑承接第三方路径。对计划在 vLLM 中集成厂商内核 (oneDNN、ACL 等) 的工程师是很好的参考模板。

功能与动机

PR body 明确说明动机: 在 AMD Zen CPU 上将 MoE 推理路由到 zentorch 融合内核, 使其优先于 AMX grouped-GEMM / OneDNN / per-expert PyTorch fallback 执行, 以利用 zentorch 针对 Zen 架构优化的融合算子提升 BF16 MoE 性能。同时补上 gelu_tanh 激活支持 (如 Gemma 4) 作为 native CPU 路径的兜底。背后是 AMD 为 vLLM CPU 后端引入厂商优化库的整体路线 (对应 AMD Zen CPU CI 的 RFC issue #44421)。

实现拆解

1. 能力检测网关 (`vllm/model_executor/kernels/linear/zentorch_utils.py`): 新增 `_ZENTORCH_MOE_ACTIVATIONS` 白名单、`_moe_activation_to_str` 归一化辅助函数, 以及两级检测函数——配置级 `is_zentorch_moe_config_supported(moe_config)` 供 `is_supported_config` 放行 zentorch 可用的配置 (绕过 grouped-gemm 对齐检查), 层级 `is_zentorch_moe_supported(layer)` 在权重加载阶段探测 op 注册、`is_act_and_mul`、`activation` 属性与白名单。所有 skip 原因走 `logging.debug`, 避免启动日志噪音 (review 中 `tlrmchlsmith` 的硬性要求)。
2. 专家层路由接入 (`vllm/model_executor/layers/fused_moe/experts/cpu_moe.py`):
CPUUnquantizedExperts.__init__ 新增 `_use_zentorch = False` 标志;
`process_weights_after_loading` 中先调用 `is_zentorch_moe_supported(layer)`, 命中则直接 return, 跳过 `_pad_moe_intermediate` 与 `cpu_prepack_moe_weight` 权重 prepack (zentorch 期望标准 [E, ...] 布局); `apply` 中若 `_use_zentorch` 为真则预分配输出并调用 `torch.ops.zentorch.zentorch_fused_moe`, 否则走既有 `cpu_fused_moe`。
`is_supported_config` 在基类、X86CPUUnquantizedExperts、ArmCPUUnquantizedExperts 三处均加入 `is_zentorch_moe_config_supported` 提前放行

逻辑。

3. 测试配套 (tests/kernels/moe/test_zen_cpu_fused_moe.py) : 新建独立测试文件, 模块级 skip 非 Zen CPU 或 op 不可用的环境 (避免在 Intel CI 空跑, 回应 fadara01 的质疑)。覆盖四类场景: 四种激活下的 dispatch 选择与权重不被 prepack、非对齐 intermediate 配置被 is_zentorch_moe_config_supported 放行、端到端 forward 与参考实现 ref_fused_moe 数值对齐 (含 bias 开关)、缺失 activation 属性时回退。
4. 依赖 pin 升级 (setup.py) : zen extra 的 zentorch 从 2.11.0.0 提升到 2.13.0.0, 确保 pip install -e .[zen] 拉取的版本包含 zentorch_fused_moe op。

关键文件:

- vllm/model_executor/kernels/linear/zentorch_utils.py (模块 内核网关; 类别 source; 类型 core-logic; 符号 _moe_activation_to_str, is_zentorch_moe_config_supported, is_zentorch_moe_supported, _ZENTORCH_MOE_ACTIVATIONS) : 新增 zentorch MoE 的两级能力检测网关 (config 级与 layer 级) 和激活白名单, 是所有路由决策的入口, 也是 review 中日志与过度工程化讨论的焦点。
- vllm/model_executor/layers/fused_moe/experts/cpu_moe.py (模块 专家内核; 类别 source; 类型 core-logic; 符号 CPUUnquantizedExperts.process_weights_after_loading, CPUUnquantizedExperts.apply, CPUUnquantizedExperts.is_supported_config, X86CPUUnquantizedExperts.is_supported_config) : CPUUnquantizedExperts 中接入 zentorch 快速路径: 权重加载阶段检测并跳过 prepack, apply 阶段调用 zentorch_fused_moe, is_supported_config 为 zentorch 配置放行非对齐形状。
- tests/kernels/moe/test_zen_cpu_fused_moe.py (模块 单元测试; 类别 test; 类型 test-coverage; 符号 _make_layer_and_config, test_zen_cpu_fused_moe_dispatches_to_zentorch, test_zen_cpu_fused_moe_config_supported_unaligned_intermediate, test_zen_cpu_fused_moe_forward) : 新增 Zen CPU 专属测试: 验证 dispatch 选择、forward 数值对齐、非对齐 intermediate 支持与缺失 activation 时的回退, 模块级 skip 避免在非 Zen 平台空转。
- setup.py (模块 构建配置; 类别 source; 类型 configuration) : zen extra 的 zentorch pin 从 2.11.0.0 提升到 2.13.0.0, 确保安装的 zentorch 版本提供 zentorch_fused_moe op。

关键符号: is_zentorch_moe_supported, is_zentorch_moe_config_supported, _moe_activation_to_str, CPUUnquantizedExperts.process_weights_after_loading, CPUUnquantizedExperts.apply, CPUUnquantizedExperts.is_supported_config

关键源码片段

vllm/model_executor/kernels/linear/zentorch_utils.py

新增 zentorch MoE 的两级能力检测网关 (config 级与 layer 级) 和激活白名单, 是所有路由决策的入口, 也是 review 中日志与过度工程化讨论的焦点。

```
# SPDX-License-Identifier: Apache-2.0
```

```
"""Gates zentorch CPU linear dispatch on platform/op availability."""
```

```

from __future__ import annotations

import logging
from typing import TYPE_CHECKING

import torch

from vllm.platforms import current_platform

if TYPE_CHECKING:
    from vllm.model_executor.layers.fused_moe.config import FusedMoEConfig

# zentorch 融合 MoE op 支持的激活函数白名单
_ZENTORCH_MOE_ACTIVATIONS = frozenset({"gelu", "gelu_tanh", "silu", "swigluoai"})

def _moe_activation_to_str(activation: object) -> str:
    """将激活函数归一化为小写字符串（兼容枚举与普通对象）。"""
    if hasattr(activation, "value"):
        return str(activation.value).lower()
    return str(activation).lower()

def is_zentorch_moe_config_supported(moe_config: FusedMoEConfig) -> bool:
    # 配置级门控：供 is_supported_config 使用，让 zentorch 可用的配置
    # 绕过 grouped-gemm 的对齐检查（例如非对齐 intermediate size）。
    if not has_zentorch_op(["zentorch_fused_moe"]):
        return False
    if not moe_config.is_act_and_mul:
        return False
    act = _moe_activation_to_str(moe_config.activation)
    return act in _ZENTORCH_MOE_ACTIVATIONS

def is_zentorch_moe_supported(layer: torch.nn.Module) -> bool:
    # 层级门控：在 process_weights_after_loading 阶段探测，
    # 只有 Zen CPU + op 已注册 + gated MLP + 激活函数在白名单内才启用。
    if not has_zentorch_op(["zentorch_fused_moe"]):
        logging.debug(
            "Skipping zentorch fused-MoE: not a Zen CPU or "
            "zentorch not loaded; using default MoE."
        )
        return False
    moe_config = getattr(layer, "moe_config", None)
    if moe_config is not None and not moe_config.is_act_and_mul:
        logging.debug(
            "Skipping zentorch fused-MoE: layer is not a gated "
            "(act-and-mul, e.g. SwiGLU) MLP, the only structure supported."
        )

```

```

        return False
activation = getattr(layer, "activation", None)
if activation is None:
    logging.debug(
        "Skipping zentorch fused-MoE: layer has no 'activation' "
        "attribute, so the activation can't be verified."
    )
    return False
act = str(getattr(activation, "value", activation)).lower()
if act not in _ZENTORCH_MOE_ACTIVATIONS:
    logging.debug(
        "Skipping zentorch fused-MoE: activation %r unsupported (supported: %s).",
        act,
        _ZENTORCH_MOE_ACTIVATIONS,
    )
    return False
return True

```

vllm/model_executor/layers/fused_moe/experts/cpu_moe.py

CPUUnquantizedExperts 中接入 zentorch 快速路径：权重加载阶段检测并跳过 prepack, apply 阶段调用 zentorch_fused_moe, is_supported_config 为 zentorch 配置放行非对齐形状。

```

def process_weights_after_loading(self, layer: torch.nn.Module) -> None:
    self.use_grouped_topk = layer.use_grouped_topk
    self.renormalize = layer.renormalize
    self.scoring_func = layer.scoring_func
    self.custom_routing_function = layer.custom_routing_function
    # 关键路由决策点：Zen CPU 上 zentorch 可用时走融合 op 快速路径，
    # 直接 return 跳过中间维度 padding 与权重 prepack，
    # 因为 zentorch 期望标准 [E, ...] 布局的原始权重。
    self._use_zentorch = is_zentorch_moe_supported(layer)
    if self._use_zentorch:
        return
    self._pad_moe_intermediate(layer)
    replace_parameter(
        layer, "w13_weight", cpu_prepack_moe_weight(layer.w13_weight, self.isa)
    )
    replace_parameter(
        layer, "w2_weight", cpu_prepack_moe_weight(layer.w2_weight, self.isa)
    )

# ... apply() 内部，位于 select_experts 与 router 权重应用之后 ...
if self._use_zentorch:
    # zentorch 融合 MoE op：输出由调用方预先分配，
    # 激活函数以小写字符串传入，全部在单个算子内完成
    # GEMM + 激活 + 权重加权，避免逐 expert 派发开销。
    output = torch.empty_like(hidden_states)
    torch.ops.zentorch.zentorch_fused_moe(

```

```

        output,
        hidden_states,
        w1,
        w2,
        self.w1_bias,
        self.w2_bias,
        topk_weights,
        topk_ids,
        apply_router_weight_on_input,
        str(activation.value).lower(),
    )
    return output
return cpu_fused_moe( # 既有 fallback: AMX/NEON/OneDNN 路径
    hidden_states,
    w1,
    w2,
    self.w1_bias,
    self.w2_bias,
    topk_weights,
    topk_ids,
    activation.value,
    self.isa,
    apply_router_weight_on_input,
)

```

评论区精华

Review 的核心交锋集中在三处：一是 fadara01 要求暂停 PR 等待 #50133 (CPU MoE 重构) 合并，作者 rebase 后将实现从旧 `CPUFusedMOE.forward_zentorch` 移植进新的 `CPUUnquantizedExperts`；二是 tlrnchlsmth 两次 CHANGES_REQUESTED，要求把启动日志降为 debug、移除无调试价值的 info_once 日志，并指出 `_moe_activation_to_str` 的 enum 兼容处理过度工程 (layer.activation 总是 MoEActivation 枚举)，作者逐一简化；三是 fadara01 对测试策略的质疑——mock zentorch op 在非 Zen 平台浪费 CI 周期、zen 专属测试不应在 Intel CI 运行，最终采纳其建议：拆出独立测试文件、删除 mock、模块级 skip。AndreasKaratzas 则从架构角度提醒“平台特定例外不应散落在通用 CPU 类中”，作者以集中 gate 函数收敛了平台判断。

- 等待 CPU MoE 重构 PR #50133 合并 (other): 等待并 rebase，最终实现落在重构后的 `CPUUnquantizedExperts` 上，重构 PR 先行合并。
- 启动日志噪音需要保持干净 (style): 改为 logging.debug，并补充更清晰的 skip 原因说明；移除冗余日志。
- 封装层 activation 归一化是否过度工程 (design): 函数被简化保留（仍用于 config 级检查），移除了 MoEActivation 本地 import 等防御逻辑。
- zen 测试应独立成文件并门控平台 (testing): 新建独立测试文件，模块级 skip 非 Zen CPU 或 op 不可用场景。

- mock zentorch op 的测试是否浪费 CI (testing): 删除 mock fixture, 改为依赖真实 zentorch op 的模块级 skip。
- 平台特定例外不应散落在通用 CPU 类中 (design): 以 zentorch_utils.py 的集中门控函数替代散落的平台判断, 获 reviewers 认可。

风险与影响

- 风险:
 1. 静默降级风险: is_zentorch_moe_supported 所有 skip 原因均为 DEBUG 级别, 若用户安装的 zentorch 版本缺少 zentorch_fused_moe op, 系统会无提示地回退到旧路径, 性能差异可能被误认为回归。
 2. 第三方 op 契约耦合: cpu_moe.py 的 apply 按位置参数调用 torch.ops.zentorch.zentorch_fused_moe, 输出缓冲由 vLLM 预分配; setup.py 将 pin 硬升级到 2.13.0.0, 若上游 op 签名或语义变化, 会导致运行时崩溃或数值偏差。
 3. 权重布局契约变化: zentorch 路径下 process_weights_after_loading 不再 prepack 权重, _use_zentorch 为真时权重保持原始布局; 任何依赖 prepack 后布局的下游逻辑 (如序列化、DP 权重传输) 都可能受影响, 需确保该标志在所有路径上一致。
 4. 测试覆盖缺口: 新测试在非 Zen CPU 平台全部 skip, 且 AMD Zen CPU CI 尚未落地 (#44421 仍在 RFC), 回归主要依赖 AMD 侧人工验证, Intel/ARM CI 无法发现该路径的破坏。
 - 影响: 用户侧: AMD Zen CPU 上 BF16 非量化 MoE 的默认推理路径改变, 预期获得融合内核带来的性能提升; Intel/ARM CPU、GPU、XPU 用户无任何行为变化。
 - 系统侧: CPUUnquantizedExperts 的生命周期语义发生变化——是否 prepack 权重取决于平台与 zentorch 可用性, 权重布局契约随运行环境不同而不同。
 - 团队侧: 确立了“能力检测 + 优雅降级”接入厂商优化内核的模式, 后续 zentorch 线性层、rms_norm 等算子可复用同一套 has_zentorch_op 门控; 同时该 PR 的合并也明确了厂商内核接入需要配套独立测试与平台门控的评审预期。
- 风险标记: 第三方内核依赖, 平台专属行为分支, 非 Zen CI 无测试覆盖, 静默降级路径

关联脉络

- PR #50133 Remove cpu_fused_moe.py (CPU MoE refactor): 本 PR 在 review 中被 fadara01 要求等待该重构合并, 随后 rebase 并把 zentorch 路径移植进重构后的 CPUUnquantizedExperts API (commit 消息明确引用 #50133)。
- PR #43344 Add gelu_tanh support to CPU fused MoE: fadara01 在 review 中指出 _gelu_tanh_and_mul 的改动已在该 PR 中实现, 本 PR 保留该函数以支持 zentorch 不可用时的 GeluTanh 路径。