

PR #43554 完整报告

vllm-project/vllm

[Kernel] Remove NormGateLinear

合并时间: 2026-05-25 17:49

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/43554>

执行摘要

- 一句话: 删除 NormGateLinear 及相关融合内核代码
- 推荐动作: 此 PR 主要是清理操作, 值得快速审阅。关注 NVIDIA 和 AMD 模型变更的一致性, 并留意 AMD 平台后续是否可融合 norm 到 MHC 以获得性能收益。

功能与动机

在 PR #43474 中, norm 融合已集成到 mhc_pre 内核, 因此 NormGateLinear 不再是必需。此 PR 旨在移除遗留代码以减少维护负担, 并自动为 DeepSeek V3 路由 GEMM 启用 PDL 支持 (通过移除相关限制)。

实现拆解

1. 删除 Python 封装层: 移除 `vllm/model_executor/layers/fused_moe/router/norm_gate_linear.py`, 包括自定义算子 `dsv4_pro_norm_gate` 和类 `NormGateLinear`。该层原是为了在 fused kernel 和 unfused fallback 之间分发的。
2. 移除 CUDA 融合内核: 删除 `csrc/moe/dsv4_norm_router_gemm.h`、`dsv4_norm_router_gemm_kernel.cu` 和 `dsv4_norm_router_gemm_entry.cu`, 以及 `moe_ops.h` 和 `dsv3_router_gemm_utils.h` 中的相关声明。这些是单内核融合 RMSNorm + GEMV 的实现。
3. 更新自定义操作绑定: 从 `vllm/_custom_ops.py` 中删除 `dsv4_norm_router_gemm` 函数, 它之前是 Python 端入口。
4. 调整模型实现: 在 `vllm/models/deepseek_v4/amd/model.py` 和 `nvidia/model.py` 中, 将 `self.norm_gate = NormGateLinear(...)` 替换为显式的 `self.gate = GateLinear(...)` 和 `self.ffn_norm = RMSNorm(...)`, 并更新 FusedMoE 初始化参数以使用 `gate` 属性。AMD 的 `_forward_cuda` 和 `_forward_rocm` 中删除注释, 显式调用 `self.ffn_norm(x)`。
5. 更新辅助内核: 在 `vllm/model_executor/kernels/mhc/tilelang.py` 中为 MHC 内核的 norm 融合添加默认 epsilon 值, 以与 fused 路径兼容。
6. 移除基准测试: 删除 `benchmarks/kernels/benchmark_norm_router_gemm.py`, 该文件不再需要因为融合内核已移除。

关键文件:

- `vllm/model_executor/layers/fused_moe/router/norm_gate_linear.py` (模块 MoE 路由层; 类别 source; 类型 deletion; 符号 `_dsv4_pro_norm_gate`,

`_dsv4_pro_norm_gate_fake`, `NormGateLinear`, `init`)：核心被删除文件，包含 `NormGateLinear` 类和 `fused` 分派逻辑，是此 PR 的主要目标。

- `vllm/models/deepseek_v4/amd/model.py`（模块 AMD 模型；类别 source；类型 data-contract）：AMD 模型适配变更量最大，需要用 `GateLinear` 替换 `NormGateLinear`，并调整 `FusedMoE` 初始化及前向传播中的显式 `norm` 调用。
- `benchmarks/kernels/benchmark_norm_router_gemm.py`（模块 性能基准；类别 source；类型 deletion；符号 `unfused_norm_router_gemm`, `fused_norm_router_gemm`, `_make_inputs`, `calculate_diff`）：删除整个基准测试文件，因为它专门测试被移除的融合内核。

关键符号：`dsv4_norm_router_gemm`, `NormGateLinear.forward`, `_dsv4_pro_norm_gate`, `_dsv4_pro_norm_gate_fake`, `GateLinear`

关键源码片段

`vllm/model_executor/layers/fused_moe/router/norm_gate_linear.py`

核心被删除文件，包含 `NormGateLinear` 类和 `fused` 分派逻辑，是此 PR 的主要目标。

```
# 已删除文件: norm_gate_linear.py
# NormGateLinear 原本封装了 RMSNorm + GateLinear 的融合或回退路径,
# 但现在融合已被 mhc_pre 内核取代, 故删除此层。

class NormGateLinear(nn.Module):
    """RMSNorm + GateLinear, fused on DSV4-Pro only."""

    def __init__(self, hidden_size, num_experts, rms_eps=1e-6, params_dtype=None, prefix=""):
        super().__init__()
        self.norm = RMSNorm(hidden_size, eps=rms_eps, dtype=params_dtype)
        self.gate = GateLinear(hidden_size, num_experts, bias=False,
                                out_dtype=torch.float32, params_dtype=params_dtype,
                                prefix=f"{prefix}.gate" if prefix else "gate")
        self._fused_kernel_supported = (
            hidden_size == 7168 and num_experts == 384
            and self.gate.allow_dsv3_router_gemm
        )

    def forward(self, x):
        if self._fused_kernel_supported:
            # fused 内核路径: 单内核完成 RMSNorm 和 GEMV (现由 mhc_pre 替代)
            return torch.ops.vllm.dsv4_pro_norm_gate(
                x, self.norm.weight, self.gate.weight, self.rms_eps
            )
        # 非 Pro 回退: 显式 norm + gate
        normed_x = self.norm(x)
        logits, _ = self.gate(normed_x)
        return normed_x, logits
```

`vllm/models/deepseek_v4/amd/model.py`

AMD 模型适配变更量最大，需要用 GateLinear 替换 NormGateLinear，并调整 FusedMoE 初始化及前向传播中的显式 norm 调用。

```
# 更新后的 AMD 模型初始化 (head 版本) # 使用显式的 GateLinear 和 RMSNorm 替代 NormGateLinear
class DeepseekV4DecoderLayer(nn.Module):
    def __init__(self, config, vllm_config, prefix):
        # ... # 原为 self.norm_gate = NormGateLinear(...)
        self.gate = GateLinear(
            input_size=config.hidden_size,
            output_size=config.n_routed_experts,
            bias=False,
            out_dtype=torch.float32,
            prefix=f"{prefix}.gate",
        )
        self.gate.e_score_correction_bias = None
        self.gate.tid2eid = None # 注意: self.ffn_norm 已在外层定义 (如 RMSNorm), 此处不再重复创建
        self.experts = FusedMoE(
            shared_experts=self.shared_experts,
            gate=self.gate, # 直接传入 gate
            num_experts=config.n_routed_experts,
            # ...
            e_score_correction_bias=self.gate.e_score_correction_bias,
            hash_indices_table=self.gate.tid2eid,
        )
        # 同时前向传播中取消注释, 显式调用 x = self.ffn_norm(x) 以执行 RMSNorm。
```

评论区精华

机器人审查者指出 AMD 实现中显式调用 `self.ffn_norm(x)` 是错失的优化机会，应融合到 MHC 操作中，因为 tilelang 内核已支持融合 norm。作者回复不确定 AMD 是否能使用该内核，未采取行动。此外，zyongye 批准了该 PR。

- AMD 平台 norm 融合优化 (performance): 作者回复 'I am not sure if AMD can utilize this kernel', 未做进一步修改。当前状态未解决。
- PR 整体批准 (other): 批准，无未解决项。

风险与影响

- 风险：功能风险极低，因为 fused 内核在 #43474 中已被替代；性能风险：AMD 平台因未利用 tilelang 融合 norm 而存在额外内核启动开销，但原本 fallback 路径也存在类似开销；构建风险：已清理 CUDA 引用，构建系统无误。
- 影响：用户：DeepSeek V4 模型推理结果不变，但推理性能可能因自动 PDL 支持而略有提升（主要影响 NVIDIA 平台）。系统：减少约 800 行代码，降低维护成本。团队：需要确保所有引用 NormGateLinear 的地方都已迁移，已检查 AMD 和 NVIDIA 模型。
- 风险标记：AMD 平台优化潜力未兑现

关联脉络

- PR #43474 [Kernel] Add mhc_pre_big_fuse_with_norm_tilelang: 此 PR 引入了融合 norm 的 mhc_pre 内核，使 NormGateLinear 变为冗余，是本 PR 清理的前提。