

PR #48476 完整报告

vllm-project/vllm

[XPU] Support MXFP8 linear weights for INC DeepSeek V4 model

合并时间: 2026-08-06 13:17

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/48476>

执行摘要

- 一句话: XPU 支持 MXFP8 线性权重, 优化 DeepSeek V4 输出投影
- 推荐动作: 值得精读。核心亮点是零额外内存的布局兼容方案: 不复制 scale, 只用 .t() 视图同时满足 checkpoint 规范布局与 oneDNN 连续布局要求; `_prepare_bmm_params` 的预计算也体现了批 GEMM 场景下减少逐调用重排开销的思路。建议结合 #49596 的讨论理解“公共算子布局假设”带来的维护困境, 并跟进 vpirogov 所要求的 oneDNN reproducer, 推动上游能力补齐。

功能与动机

PR body 明确说明目标是让 DeepSeek V4 XPU 输出投影路径能够高效消费 compressed-tensors MXFP8 权重。此前 XPU 的 MXFP8 kernel 在加载时直接转置权重与 scale, 破坏了 checkpoint 的规范布局, 导致其他公共算子 (如 MLA 反量化路径) 出现 shape mismatch; 本 PR 通过保留原始布局并运行时视图转置来解决兼容性问题, 同时为 BMM 层预计算批 GEMM 参数以消除逐调用转置开销。

实现拆解

该 PR 只修改一个文件 `vllm/model_executor/kernels/linear/mx_fp8/xpu.py`, 实现分三步:

1. 权重与 scale 布局策略调整: 在 `process_weights_after_loading` 中, 不再把 `layer.weight` 转置替换, 也不再把 `layer.weight_scale` 直接替换为连续 `[K//32, N]` 缓冲; 而是将 checkpoint 的 `[N, K//32]` scale 先转置并 `contiguous()` 得到 `scale_kn`, 再以 `scale_kn.t()` 视图存回 `layer.weight_scale`。这样外部消费者看到的 shape 仍是 checkpoint 的 `[N, K//32]`, 而 `apply_weights` 里 `.t()` 就能零成本拿回 oneDNN 需要的连续 `[K//32, N]` 布局。
2. 新增 `_prepare_bmm_params`: 当 layer 带有 `is_bmm=True` 时 (如 DeepSeek V4 的 `wo_a`), 根据 `bmm_batch_size` 将 scale `[K//32, N_total]` reshape 为 `[G, K//32, N_per_group]` 并重排为连续 `bmm_scale`, 同时把 weight `[N_total, K]` reshape 为连续 `bmm_weight[G, K, N_per_group]`, 避免每次前向都做 `permute().contiguous()`。
3. `apply_weights` 运行时视图转置: 调用 `torch.ops._xpu_C.fp8_gemm` 时传 `layer.weight.t()` 与 `layer.weight_scale.t()`, 两个视图都是零拷贝; 同时保持 `layer.weight` 原始 `[N, K]` 不动, 避免影响其他假设原始形状的算子。

4. 测试与配套：本 PR 没有新增单元测试，模型侧验证为 GSM8K 8x B70 精度 0.948；CI 仅由维护者触发 Intel 相关流水线。

关键文件：

- `vllm/model_executor/kernels/linear/mx_fp8/xpu.py`（模块 MXFP8 内核；类别 source；类型 data-contract；符号 `_prepare_bmm_params`, `process_weights_after_loading`, `apply_weights`）：唯一改动文件，承载了布局策略调整、BMM 参数预计算与运行时视图转置三处核心逻辑，直接决定 XPU MXFP8 权重能否被 oneDNN 高效消费且不破坏其他算子的布局假设。

关键符号：`_prepare_bmm_params`, `process_weights_after_loading`, `apply_weights`

关键源码片段

`vllm/model_executor/kernels/linear/mx_fp8/xpu.py`

唯一改动文件，承载了布局策略调整、BMM 参数预计算与运行时视图转置三处核心逻辑，直接决定 XPU MXFP8 权重能否被 oneDNN 高效消费且不破坏其他算子的布局假设。

```
# vllm/model_executor/kernels/linear/mx_fp8/xpu.py
# XPUMxP8LinearKernel: XPU 上 MXFP8 W8A8 GEMM 内核。

def process_weights_after_loading(self, layer: torch.nn.Module) -> None:
    # 检查点中 scale 形状为 [N, K//32]（每 32 个元素一个 E8M0 scale）。
    # oneDNN 的 fp8_gemm 要求连续 [K//32, N] 布局。
    # 这里把 transposed contiguous 缓冲存成 .t() 视图，使得：
    # - 其他消费者（如反量化路径）仍然看到 checkpoint 的形状 [N, K//32]；
    # - apply_weights 里只需 .t() 即可零成本拿回 oneDNN 需要的连续布局。
    weight_scale = layer.weight_scale.view(torch.float8_e8m0fnu)
    scale_kn = weight_scale.data.t().contiguous()
    replace_parameter(layer, "weight_scale", scale_kn.t())

    # 对 BMM 类层（如 DeepSeek V4 的 wo_a）预计算批 GEMM 权重与 scale
    if getattr(layer, "is_bmm", False):
        self._prepare_bmm_params(layer, scale_kn)

def _prepare_bmm_params(
    self, layer: torch.nn.Module, scale_kn: torch.Tensor
) -> None:
    """预计算批处理权重与 scale，供分组 fp8_bmm 使用（例如 wo_a）。

    将 scale [K//32, N_total] 切分为 [G, K//32, N_per_group]，
    将 weight [N_total, K] 整理为连续 [G, K, N_per_group]。
    """
    batch = layer.bmm_batch_size
    k_blocks, n_blocks = scale_kn.shape
    layer.bmm_scale = (
        scale_kn.reshape(k_blocks, batch, n_blocks // batch)
        .permute(1, 0, 2)
```

```

        .contiguous()
    )
    w = layer.weight
    n_total, k = w.shape
    layer.bmm_weight = (
        w.reshape(batch, n_total // batch, k).permute(0, 2, 1).contiguous()
    )

def apply_weights(
    self,
    layer: torch.nn.Module,
    x: torch.Tensor,
    bias: torch.Tensor | None = None,
) -> torch.Tensor:
    out_dtype = x.dtype
    x_fp8, x_scale = quant_mxfp8(x)
    # weight 保存为 [N, K]; .t() 得到 [K, N] 视图，不拷贝。
    # scale 保存为 [N, K//32] 视图; .t() 恢复 oneDNN 期望的连续 [K//32, N] 缓冲。
    return torch.ops._xpu_C.fp8_gemm(
        x_fp8,
        layer.weight.t(),
        out_dtype,
        x_scale,
        layer.weight_scale.t(),
        bias,
    )

```

评论区精华

审阅中围绕 [xpu.py:77](#) 的 `apply_weights` 改动展开了多轮讨论：

- zufangzhu 首先质疑为什么在 `apply_weights` 中改为 `.t()` 视图；xwu-intel 解释是为了保留原始 `weight/scale` 形状，因为其他算子（如反量化）可能假设原始布局。
- zufangzhu 进一步担心每个 Linear 额外保有一份转置 `scale` 的内存开销，xwu-intel 给出数据：TP=8 时 DeepSeek-V4-Flash 约 57.75 KiB/GPU，MXFP4-Mixed-CT-AutoRound 约 28.875 MiB/GPU；zufangzhu 指出 dense 模型所有 attention/MLP Linear 都会走此路径，开销可能更高。
- yma11 建议采用 PR #49596 的“layer attr 标识 `scale` 是否转置”方案，xwu-intel 认为公共算子众多难以逐处维护，最好由 oneDNN 直接支持这种 `scale` 布局；vpirogov 也期望一个 PyTorch 级 reproducer 说明 oneDNN 缺失点。
- 最终讨论结论：短期在 vLLM 侧用 `scale_kn.t()` 视图替换原 `scale`，不产生额外拷贝；长期推动 oneDNN 支持非连续 / 转置 `scale` 布局。
- 无关代码格式改动 (style): xwu-intel 已 revert，保持改动聚焦。
- `apply_weights` 中改为 `.t()` 视图的原因 (design): 采用运行时 `.t()` 视图方案，既满足 oneDNN 布局又不影响其他消费者。

- 额外 scale 拷贝的内存开销 (performance): 最终实现改为复用 `scale_kn.t()` 视图, 不产生额外拷贝, 内存担忧消除。
- vLLM 侧拷贝 vs 推动 oneDNN 支持 (design): 短期在 vLLM 侧用转置视图解决, 长期推动 oneDNN 支持任意 scale 布局, 交由 jikunshang 与 zufangzhu 决策。
- oneDNN 缺失能力与 reproducer (question): 本 PR 合并时 reproducer 尚未提供, 属于遗留的跨团队跟进项。

风险与影响

- 风险: 主要风险集中在布局契约与测试覆盖上:
- 布局契约变更: `process_weights_after_loading` 现在把 `layer.weight_scale` 存为 `scale_kn.t()` 视图 (非连续), 虽然 shape 保持 $[N, K/32]$, 但任何依赖底层 contiguity 或 storage offset 的算子可能出错; 必须确认所有反量化消费者都只读取 shape 而非依赖连续内存。
- BMM 路径依赖属性: `_prepare_bmm_params` 依赖 `layer.is_bmm` 与 `layer.bmm_batch_size`, 如果其他模型误设这些属性或未正确设置, 会触发路径错误; 当前仅 DeepSeek V4 使用。
- 缺少单元测试: 没有针对 kernel 布局的回归测试, 唯一验证是 GSM8K 精度; 后续重构很容易破坏现有布局约定而无人察觉。
- oneDNN 兼容性: 当前靠 `.t()` 视图规避 oneDNN 对连续 scale 的要求, 若 oneDNN 后续行为变化或未来支持非连续 scale, 此处的依赖需要同步更新。
- 影响: 影响范围限于 XPU 平台使用 MXFP8 线性权重的模型, 特别是 DeepSeek V4 及其 INC 产出的 compressed-tensors/AutoRound 检查点; 理论上所有走 `XPUMxFp8LinearKernel` 的 XPU 模型 (包括 dense 模型) 都会受到布局策略调整的影响, 但行为与原先等价。对 vLLM 团队而言, 此 PR 引入的“保留规范布局 + 运行时视图转置”模式可作为平台特定 kernel 兼容其他公共算子的参考; 与 oneDNN 团队的协作也会影响未来 XPU 量化算子设计。
- 风险标记: 量化布局契约变更, 缺少单测覆盖, 平台特定路径影响面有限, oneDNN 兼容性待验证

关联脉络

- 暂无明显关联 PR