

PR #50349 完整报告

vllm-project/vllm

[XPU] Fix FP8 block scale layout for MLA compatibility

合并时间: 2026-07-31 11:05

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/50349>

执行摘要

- 一句话: 修复 XPU FP8 block scale 布局与 MLA 兼容
- 推荐动作: 该 PR 值得精读, 尤其是对 XPU 或量化相关开发人员。关键设计决策是使用 `.t()` 视图来同时满足不同消费者对 scale 形状的期望, 避免了数据拷贝。建议阅读 `process_weights_after_loading` 和 `apply_block_scaled_mm` 的完整实现。

功能与动机

原始代码将 scale 直接存储为连续的 `[k_blocks, n_blocks]` 布局, 但 MLA 的 `process_weights_after_loading` 调用 `scaled_dequantize` 时期望 scale 形状与权重布局 `[N, K]` 匹配, 即 `[n_blocks, k_blocks]`, 导致形状不匹配和断言失败。该 PR 旨在解决此兼容性问题。

实现拆解

实现拆解

1. 调整 scale 存储布局 (`vllm/model_executor/kernels/linear/scaled_mm/xpu.py`): 在 `process_weights_after_loading` 中, 将 checkpoint 的 `[n_blocks, k_blocks]` scale 转置为连续的 `[k_blocks, n_blocks]` 缓冲区, 但通过 `replace_parameter` 存储其 `.t()` 视图, 使外部消费者 (如 MLA 的 `scaled_dequantize`) 看到 `[n_blocks, k_blocks]` 形状。
2. 计算时恢复布局: 在 `apply_block_scaled_mm` 中, 对 `Bs` 调用 `.t()` 恢复连续的 `[k_blocks, n_blocks]` 缓冲区, 以满足 oneDNN 的 `fp8_gemm` 要求, 无需额外数据拷贝。
3. 抽取 BMM 参数准备: 将原本内联在 `process_weights_after_loading` 中的 BMM 参数预处理逻辑抽取为独立的 `_prepare_bmm_params` 方法, 接收转置后的 scale (`[k_blocks, n_blocks]`), 按 batch 拆分 scale 和 weight 为 3D 张量, 供 grouped `fp8_bmm` 使用。
4. 测试与验证: 本 PR 未增加新的单元测试, 但作者在 PR body 中报告了相关测试通过情况 (`test_can_initialize_large_subset[DeepseekV3ForCausalLM]`) 和 GSM8K 评估结果 (94.92% vs 94.01%)。

关键文件:

- `vllm/model_executor/kernels/linear/scaled_mm/xpu.py` (模块内核层; 类别 `source`; 类型 `data-contract`; 符号 `_prepare_bmm_params`): 核心变更文件, 修正 scale 布局以兼容 MLA 和 oneDNN, 并抽取 BMM 参数准备逻辑。

关键符号：process_weights_after_loading, _prepare_bmm_params, apply_block_scaled_mm

关键源码片段

vllm/model_executor/kernels/linear/scaled_mm/xpu.py

核心变更文件，修正 scale 布局以兼容 MLA 和 oneDNN，并抽取 BMM 参数准备逻辑。

关键源码片段

```
# vllm/model_executor/kernels/linear/scaled_mm/xpu.py
class XPUPf8BlockScaledMMKernel(Fp8BlockScaledMMLinearKernel):
    def process_weights_after_loading(self, layer: torch.nn.Module):
        super().process_weights_after_loading(layer)
        scale_attr = (
            "weight_scale_inv" if hasattr(layer, "weight_scale_inv") else "weight_scale"
        )
        scale = getattr(layer, scale_attr)

        # Checkpoint scale is [n_blocks, k_blocks] (one value per 128x128 tile).
        # oneDNN fp8_gemm requires contiguous [k_blocks, n_blocks] layout.
        # 我们存储转置后的连续缓冲区作为 .t() 视图，使得：
        # - MLA 的 scaled_dequantize 仍能看到 [n_blocks, k_blocks] 形状
        # - apply_block_scaled_mm 通过 .t() 恢复连续缓冲区
        scale_kn = scale.data.t().contiguous() # [k_blocks, n_blocks]
        replace_parameter(layer, scale_attr, scale_kn.t()) # view: [n_blocks, k_blocks]

        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:
        """Precompute batched weight and scale for grouped fp8_bmm (e.g. wo_a)."""
        # 拆分 scale [k_blocks, n_blocks] 为 [G, k_blocks, n_blocks_per_group]
        # 以及 weight [N_total, K] 为 [G, K, N_per_group] 用于批量 GEMM。
        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
        ) # [G, K, N_per_group]

    def apply_block_scaled_mm(self, A, B, As, Bs) -> torch.Tensor:
        # B 为 [N, K]; .t() 得到 [K, N] 视图（无拷贝）。
```

```
# Bs 存储为 [n_blocks, k_blocks] 视图; .t() 恢复 oneDNN 期望的连续 [k_blocks, n_blocks]
缓冲区。
return torch.ops._xpu_C.fp8_gemm(
    A,
    B.t(),
    self.config.out_dtype,
    As,
    Bs.t(),
    torch.Tensor(),
)
```

评论区精华

本 PR 的 review 评论较少，未发现实质性的技术争议。多数 reviewer（如 xinyu-intel、xwu-intel 等）均批准了该变更。yma11 仅评论“LGTM. Thanks.”，表明变更获得认可。

- 暂无高价值评论线程

风险与影响

- 风险：## 风险分析
- 兼容性风险：变更了 scale 的存储方式，可能影响其他依赖 scale 形状的代码路径。目前测试覆盖有限，仅验证了 DeepseekV3 的初始化，未覆盖其他模型或算子。
- 性能风险：Bs.t() 操作理论上零拷贝，但需确认在 XPU 上 .t() 是否真正不产生拷贝，以及是否影响 oneDNN 的输入布局要求。
- 回归风险：BMM 参数准备逻辑被抽取为独立方法，若调用顺序或参数传递有误，可能影响 wo_a 等 BMM 层的计算。
- 影响：## 影响分析
- 用户影响：修复了 XPU 上 FP8 量化模型的初始化问题，使 DeepseekV3 等模型能够在 XPU 上正常加载。
- 系统影响：涉及 XPU 算子层的数据布局，可能影响所有使用 XPUFp8BlockScaledMMKernel 的模型。
- 团队影响：增加了代码的可读性和可维护性（抽取 _prepare_bmm_params）。
- 风险标记：缺少测试覆盖，数据布局变更，XPU 平台特定

关联脉络

- PR #50434 [XPU] [BugFix] Add deepseek_v4_fp8 to xpu supported_quantization list: 同为 XPU 平台相关，涉及 FP8 量化支持，可能与本次变更相互影响。
- PR #46516 Enable gfx1250 ROCm architecture: 涉及量化内核和平台支持，与本 PR 在量化路径上有潜在关联。