

PR #52118 完整报告

vllm-project/vllm

[XPU] [Bugfix] process ragged weights in xpu linear backend

合并时间: 2026-08-14 09:11

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/52118>

执行摘要

- 一句话: XPU 线性后端处理 ragged 权重 scale
- 推荐动作: 该 PR 值得 XPU/oneDNN 量化内核维护者精读。设计亮点在于用 gcd 选取同时整除 N 与 block_n 的分组宽度, 避免新增权重数据, 并对不支持的 ragged K 直接 fail loudly; 注释也清楚解释了 oneDNN 布局约束。建议补一个针对 N 非对齐 shape 的单元测试, 覆盖 index_select 索引正确性。

功能与动机

PR body 给出了复现命令: 在 XPU 上运行 `python examples/basic/offline_inference/generate.py --model gaunernst/DeepSeek-V2-Lite-Chat-FP8 --enforce-eager --max-model-len 2048 --trust-remote-code`, 修复前输出异常 (before/after 截图对比), 修复后正常。根因是 FP8 块量化权重 N 维度不整除 block_n=128, 而 oneDNN fp8_gemm 要求 n_blocks 整除 N。

实现拆解

1. 在 `vllm/model_executor/kernels/linear/scaled_mm/xpu.py` 的 `XPUFp8BlockScaledMMKernel.process_weights_after_loading` 中, 先读取 `self.weight_group_shape` 得到 block_n、block_k, 再读取 `layer.weight.shape` 得到 N、K。
2. 当 $N \% \text{block_n} \neq 0$ 时, 计算 $gn = \text{math.gcd}(N, \text{block_n})$, 并断言 $gn \% 16 == 0$ (oneDNN 要求分组宽度为 16 的倍数)。随后通过 `col_start` 与 `src_idx` 使用 `index_select` 将 `scale` 从 `[ceil(N/block_n), K/block_k]` 扩展为 `[N/gn, K/block_k]`, `weight` 本身不动。
3. 对 ragged K 显式拒绝: 断言 $K \% \text{block_k} == 0$, 因为 ragged K 需要同时扩展运行期激活 `scale`, 当前未处理; DeepSeek/GLM 的 K 均保持块对齐, 因此直接 fail loudly。
4. 保持既有转置与参数替换流程不变: `scale_kn = scale.data.t().contiguous()` 得到 `[k_blocks, n_blocks]` 连续布局, `replace_parameter` 存回 `.t()` 视图, `is_bmm` 分支继续调用 `_prepare_bmm_params`。
5. 配套与测试: 无新增测试文件, 仅单文件 +29/-1; 提交演进从最初实现 (8e3da56) 到注释整理 (cb14cfe), 再到补充 $gn \% 16 == 0$ 断言 (e63a140)。

关键文件:

- `vllm/model_executor/kernels/linear/scaled_mm/xpu.py` (模块 量化内核; 类别 `source`; 类型 `data-contract`; 符号 `process_weights_after_loading`): 唯一变更文件, 核心修复逻辑全部集中于此: 在 `XPUFp8BlockScaledMMKernel.process_weights_after_loading` 中处理 ragged N 的 `scale` 扩展, 并对 ragged K 显式断言。

关键符号: `process_weights_after_loading`

关键源码片段

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

唯一变更文件, 核心修复逻辑全部集中于此: 在 `XPUFp8BlockScaledMMKernel.process_weights_after_loading` 中处理 ragged N 的 `scale` 扩展, 并对 ragged K 显式断言。

```
class XPUFp8BlockScaledMMKernel(Fp8BlockScaledMMLinearKernel):
    # ... is_supported 等其它方法略去 ...

    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)

        # Ragged N (N % block_n != 0) 场景: oneDNN fp8_gemm 要求 n_blocks 能整除 N。
        # 权重本身不动, 只把 scale 的行重复扩展到更细的 N 分组 gn,
        # 其中 gn = gcd(N, block_n) 同时整除 N 与 block_n:
        # scale 从 [ceil(N/block_n), K/block_k] 变为 [N/gn, K/block_k]
        # oneDNN 只接受 gn 为 16 的倍数, 而 block_n = 128 时 gcd 必为 2 的幂,
        # 所以 gn 要么 >= 16, 要么报错。N % block_n == 0 时此分支不执行。
        block_n, block_k = self.weight_group_shape
        N, K = layer.weight.shape
        if N % block_n != 0:
            gn = math.gcd(N, block_n)
            assert gn % 16 == 0, (
                f"XPU block-scaled FP8: N ({N}) yields group width {gn}, but "
                f"oneDNN only supports multiples of 16; this weight shape is "
                f"unsupported."
            )
            col_start = torch.arange(N // gn, device=scale.device) * gn
            src_idx = torch.div(col_start, block_n, rounding_mode="floor")
            scale = scale.index_select(0, src_idx).contiguous()

        # Ragged K 需要同时扩展运行期激活 scale, 当前不处理;
        # DeepSeek/GLM 的 K 均保持块对齐, 因此直接显式报错。
        assert K % block_k == 0, (
            f"XPU block-scaled FP8 requires K ({K}) to be a multiple of the "
            f"weight block size ({block_k}); ragged-K weights are unsupported."
        )
```

```

# checkpoint 里的 scale 是 [n_blocks, k_blocks] (每个块一个值) 。
# oneDNN fp8_gemm 需要连续的 [k_blocks, n_blocks] 布局,
# 这里把转置后的连续 buffer 以 .t() 视图存回 layer, 保证:
# - MLA scaled_dequantize 仍能看到 [n_blocks, k_blocks] 形状
# - apply_block_scaled_mm 通过 .t() 取回连续 buffer
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)

```

评论区精华

没有实质性的 review 讨论。claude[bot] 因 PR 来自 fork 而自动跳过 review；维护者 jikunshang 直接 APPROVED，未留下评审意见。提交历史反映出对断言逐步加强的过程，最后专门补充了 `gn % 16 == 0` 的硬校验。

- 暂无高价值评论线程

风险与影响

- 风险:
 1. 数据契约风险: scale 的 shape 从 $\lceil N/\text{block_n} \rceil, K/\text{block_k}$ 变为 $[N/\text{gn}, K/\text{block_k}]$ ，虽然通过 .t() 保持下游逻辑可见形状，但 n_blocks 数量发生变化，需确保 MLA scaled_dequantize 与 apply_block_scaled_mm 对扩展后的 shape 语义一致。
 2. 硬断言风险: `gn % 16 == 0` 与 `K % block_k == 0` 会在不支持的权重形状下直接抛错（如 N 小于 16 或 K 不齐），若未来出现此类模型将启动失败，属于有意设计的 fail loudly。
 3. 回归风险: 没有对应的单元测试，index_select 的索引计算（col_start、torch.div）依赖 shape 假设，如果上游权重加载逻辑变化可能产生越界或错位。
 4. 影响面: 仅 XPU + block-scaled FP8 路径，不影响 CUDA/ROCm 后端。- 影响: 用户侧: XPU 上使用 FP8 块量化且 N 维度不齐的模型（如 DeepSeek-V2-Lite-Chat-FP8）从无法运行变为可正常离线推理。系统侧: 单文件局部修改，无全局配置或接口变更，不影响其他平台。团队侧: 为 Intel XPU 后端补上了一个实际可复现的兼容性缺口，但建议后续补充单元测试防止回归。- 风险标记: 缺少测试覆盖，数据契约变更，硬断言拒绝形状

关联脉络

- PR #50534 [XPU] Add tuned Mamba SSU configs for Intel Arc Pro B70: 同为 XPU 平台的 kernel 级改动，展示了 Intel GPU 后端的持续演进，可对照理解 XPU 后端的通用约束。
- PR #51793 [Quantization] Remove dead QuantizationConfig.is_mxfp4_quant: 同为 FP8/ 量化路径的清理与收敛，可对照理解 block-scaled FP8 的配置与数据契约。

- PR #48666 [Kernel] Gemma-4 FA4 FP8 Kernel: 同为 FP8 kernel 相关改动，但是面向 NVIDIA FA4 路径，可对比不同平台后端在 FP8 处理上的差异。