

PR #50787 完整报告

vllm-project/vllm

[XPU] Route block-quantized FP8 weights to the W8A8 kernel

合并时间: 2026-08-12 16:57

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/50787>

执行摘要

- 一句话: XPU 上 block 量化 FP8 权重默认改走 W8A8 内核, 修复启动崩溃
- 推荐动作: 推荐 XPU/ 量化相关工程师精读: 这是一个小型但典型的平台差异处理案例, 展示了如何在 kernel 能力不对称时通过 scheme 选择逻辑做保守的路径修正。值得关注的设计决策是“当目标 kernel 不支持某量化粒度时, 强制降级到支持该粒度的另一 kernel, 而不改变其他权重的 opt-in 语义”。同时建议在社区后续推进 #49664/#50826 的 torch backend block-wise 支持时, 将本 PR 的 XPU 专属分支纳入统一方案。

功能与动机

PR body 明确指出: 修复 #43645 之后的问题, 将 block 量化权重无条件路由到 W8A8 kernel (无需 opt-in flag `--linear-backend xpu`), 同时保持 per-tensor/per-channel 权重原有 opt-in 行为。无本 PR 时, 服务启动失败, worker 在 `load_model` 期间崩溃, 报错为 `ValueError: Failed to find a kernel that can implement the ScaledMM linear layer. XPUW8A16FP8Linear only support per-channel and per-tensor quantization`, 根因是 block 量化权重默认仍被路由到 `CompressedTensorsW8A16Fp8`, 而 `XPUW8A16FP8LinearKernel` 的 `can_implement` 显式拒绝 BLOCK 策略。

实现拆解

1. 变更入口: `vllm/model_executor/layers/quantization/compressed_tensors/compressed_tensors.py` 的 `_get_scheme_from_parts`, 该函数负责将 `compressed-tensors` 的量化格式映射到具体 kernel scheme。
2. 核心逻辑: 在 `self._is_fp8_w8a8(weight_quant, input_quant)` 的 XPU 分支内, 新增 `weight_quant_is_block_strategy` 判断, 即 `weight_quant` 存在且 `strategy == QuantizationStrategy.BLOCK` 时, 即使当前 `is_fp8_w8a8_supported` 为假也强制置为真。这样 block 量化权重会生成 `CompressedTensorsW8A8Fp8` scheme, 进而选用支持 BLOCK 策略的 `XPUFp8BlockScaledMMKernel`; per-tensor/per-channel 权重仍按原逻辑仅在 `--linear-backend xpu` 或 torch 下走 W8A8。
3. 配套 CI: `.buildkite/intel_jobs/test-intel.yaml` 的 XPU compressed tensors FP8 test 步骤追加运行 `test_compressed_tensors_fp8_block_enabled` 用例, 确保 block 量化路由在 Intel 硬件 CI 中持续覆盖; 该测试文件已在 `source_file_dependencies` 中, 无需新增测试文件。

关键文件：

- `vllm/model_executor/layers/quantization/compressed_tensors/compressed_tensors.py` (模块 量化层；类别 source；类型 core-logic；符号 `_get_scheme_from_parts`)：核心变更文件：在 `_get_scheme_from_parts` 的 XPU 分支中强制为 block 量化权重启用 W8A8 FP8 kernel，修复模型加载崩溃。
- `.buildkite/intel_jobs/test-intel.yaml` (模块 构建配置；类别 config；类型 configuration)：在 XPU compressed tensors FP8 测试步骤中追加 block 量化用例，确保新路由路径在 Intel CI 中持续被覆盖。

关键符号：`_get_scheme_from_parts`

关键源码片段

`vllm/model_executor/layers/quantization/compressed_tensors/compressed_tensors.py`

核心变更文件：在 `_get_scheme_from_parts` 的 XPU 分支中强制为 block 量化权重启用 W8A8 FP8 kernel，修复模型加载崩溃。

```
# 位于 vllm/model_executor/layers/quantization/compressed_tensors/compressed_tensors.py
# 的 _get_scheme_from_parts 中，XPU 平台 W8A8 / W8A16 FP8 内核的选择逻辑。
```

```
if self._is_fp8_w8a8(weight_quant, input_quant):
```

```
    if current_platform.is_xpu():
```

```
        # 在 XPU 上，--linear-backend xpu 或 torch 会显式启用
```

```
        # W8A8 FP8 linear kernel；否则默认回退到 W8A16。
```

```
        config = get_current_vllm_config_or_none()
```

```
        is_fp8_w8a8_supported = config is not None and (
            config.kernel_config.linear_backend in ("xpu", "torch")
        )
```

```
        # 判断权重量化是否为 block 粒度（例如 FP8-block 模型）
```

```
        weight_quant_is_block_strategy = (
```

```
            weight_quant
```

```
            and weight_quant.strategy == QuantizationStrategy.BLOCK
```

```
        )
```

```
        if weight_quant_is_block_strategy and not is_fp8_w8a8_supported:
```

```
            # XPU 的 W8A16 kernel 不支持 BLOCK 策略，
```

```
            # 因此 block 量化权重无条件改走 W8A8 kernel。
```

```
            is_fp8_w8a8_supported = True
```

```
    else:
```

```
        # 非 XPU 平台按 GPU capability 判断是否支持 W8A8
```

```
        is_fp8_w8a8_supported = self._check_scheme_supported(
```

```
            CompressedTensorsW8A8Fp8.get_min_capability(), error=False
```

```
        )
```

```
    if is_fp8_w8a8_supported:
```

```
        # 走支持 BLOCK 策略的 W8A8 scheme (XPU 上为 XPUPp8BlockScaledMMKernel)
```

```
        return CompressedTensorsW8A8Fp8(
```

```

        weight_quant=weight_quant,
        is_static_input_scheme=(input_quant and not input_quant.dynamic),
    )
else:
    # input_quant 在转换模型中会保留，但推理阶段会被忽略
    return CompressedTensorsW8A16Fp8(
        weight_quant=weight_quant,
        is_static_input_scheme=not input_quant.dynamic,
    )

```

评论区精华

@zufangzhu 提醒作者关注 #49664: torch backend 已支持 w8a8 per-tensor 和 per-channel/per-token，且 #50826 合并后会启用 block-wise kernel，建议考虑通用性；作者回应已加入“w8a8 未启用且策略为 block_quant 时置为 true”的逻辑。mergify[bot] 两次报告 merge conflict，要求 rebase，最终解决并触发 CI (Buildkite #83523)。jikunshang 先 cc @zufangzhu 请求量化方向确认，随后直接批准。

- torch backend 对 w8a8/block 的支持与方案对齐 (design): @chaojun-zhang 回应已加入逻辑：若 w8a8 未启用且策略为 block_quant，则强制置为 true。
- merge conflict 处理 (other): 作者完成 rebase 后由 jikunshang 触发 /ci run 并通过。
- 评审流转与批准 (other): 已批准并合并。

风险与影响

- 风险：
 1. 核心加载路径变更：scheme 选择发生在模型加载的 create_weights 阶段，影响所有 XPU 上使用 compressed-tensors FP8 block 量化权重的模型，若 XPUFp8BlockScaledMMKernel 对某些 activation 量化组合（如动态 per-token）存在精度问题，会直接影响推理结果。
 2. 平台限定分支：改动仅在 current_platform.is_xpu() 为真时生效，NVIDIA/AMD 分支逻辑完全不变，回归风险低。
 3. 测试配套薄弱：仅通过 CI 配置启用已有测试用例，没有新增针对 BLOCK 策略 scheme 选择的单元测试，覆盖依赖 Intel 硬件 CI 的实机运行。
 4. 未来演进风险：若 XPU 的 W8A16 kernel 后续补齐 BLOCK 支持，该强制逻辑需要回收，否则会一直绕开 W8A16 路径。- 影响：影响范围限定在 XPU (Intel GPU) 平台：此前无法启动的 block 量化 FP8 模型（如 Qwen3-32B-FP8-block）现在无需任何 opt-in 参数即可正常加载和推理；per-tensor/per-channel FP8 权重行为不变。对 NVIDIA/AMD 等平台无影响。改动量小 (+11/-1)，单 commit，合并风险低。对团队而言，该修复补齐了 XPU 多粒度量化支持的最后一个默认路径缺口。- 风险标记：核心加载路径变更，平台限定分支 (XPU)，测试配套薄弱（仅 CI 启用已有测试）

关联脉络

- PR #43645 [XPU] Add W8A8 FP8 linear kernel with multi-granularity quant support:
本 PR 的直接前驱：引入 XPU W8A8/W8A16 内核选择逻辑，本 PR 修复其中 block 量化权重默认路由失败的问题。