

PR #50840 完整报告

vllm-project/vllm

[XPU] Route AWQ linear through choose_mp_linear_kernel

合并时间: 2026-08-06 13:50

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/50840>

执行摘要

- 一句话: XPU AWQ 统一走 choose_mp_linear_kernel, 删除专用实现
- 推荐动作: 值得精读, 尤其是 choose_mp_linear_kernel 统一平台内核选择的设计思路: 通过将 AWQ 权重统一转换为 GPTQ 格式, 让一个 LinearMethod 服务多个平台的 WNA16 内核, 是消除平台分叉的典型模式。关注 process_weights_after_loading 中权重转换的契约, 以及未来 XPU 内核扩展时是否需要重新引入平台判断。

功能与动机

PR 正文指出, #42727 已统一 AWQ 配置并修复 CPU, 但 XPU 仍保留专用 AutoAWQXPULinearMethod, 形成平台分歧; 同时 XPU 一旦回落到 AutoAWQLinearMethod, 其 apply() 会调用 CUDA-only 的 ops.awq_dequantize, 这正是 issue #41469 的根因。因此需要让 XPU 走与 CPU 一致的 choose_mp_linear_kernel 路径, 消除分歧并避免对 CUDA 专属算子的依赖。

实现拆解

1. 统一内核选择入口: 在 vllm/model_executor/layers/quantization/auto_awq.py 的 AutoAWQConfig.get_quant_method 中, 将原先独立的 is_xpu() 专用分支与 is_cpu() 分支合并为 is_cpu() or is_xpu(), 统一返回 AutoAWQMarlinLinearMethod, 由 choose_mp_linear_kernel 选择内核 (CPU 为 CPUWNA16LinearKernel, XPU 为 XPUwNa16LinearKernel)。注释同步更新, 说明 AWQ 权重会在 process_weights_after_loading 中转换为标准 GPTQ 格式后再交给内核。
2. 扩展 Marlin 验证豁免: 在 AutoAWQMarlinLinearMethod.__init__ 中, 将 verify_marlin_supported 的跳过条件从仅 CPU 扩展为 CPU/XPU, 因为两个平台均使用专用 WNA16 内核, 不经过 Marlin 校验路径。
3. 删除专用实现: 移除 AutoAWQXPULinearMethod 整个类 (process_weights_after_loading、_get_group_size、apply 及对 vllm_xpu_kernels 的依赖), 净删除约 70 行代码, 消除并行维护的权重转换契约。
4. CI 配套: 在 .buildkite/intel_jobs/quantization.yaml 的 Intel XPU Quantization job 命令中追加 tests/quantization/test_auto_awq.py, 使 XPU AWQ 路径 (包括 Qwen2-1.5B-Instruct-AWQ 端到端冒烟测试) 进入 CI 覆盖。此外, 提交历史显示两次 merge main (4746ea6、7de4961), 说明分支与主干保持同步。

关键文件：

- vllm/model_executor/layers/quantization/auto_awq.py（模块 量化层；类别 source；类型 core-logic；符号 AutoAWQConfig.get_quant_method, AutoAWQMarlinLinearMethod.init, AutoAWQXPULinearMethod, AutoAWQXPULinearMethod.process_weights_after_loading）：核心变更文件：统一 XPU 与 CPU 的 AWQ 内核选择路径，扩展 Marlin 验证豁免，并删除 AutoAWQXPULinearMethod 整体实现。
- .buildkite/intel_jobs/quantization.yaml（模块 CI 配置；类别 config；类型 configuration）：在 Intel XPU 量化 CI 中补充 AWQ 测试，覆盖新的内核路由路径，防止回归。

关键符号：AutoAWQConfig.get_quant_method, AutoAWQMarlinLinearMethod.init, AutoAWQXPULinearMethod.apply（移除），AutoAWQXPULinearMethod.process_weights_after_loading（移除），AutoAWQXPULinearMethod._get_group_size（移除）

关键源码片段

vllm/model_executor/layers/quantization/auto_awq.py

核心变更文件：统一 XPU 与 CPU 的 AWQ 内核选择路径，扩展 Marlin 验证豁免，并删除 AutoAWQXPULinearMethod 整体实现。

```
# vllm/model_executor/layers/quantization/auto_awq.py
# 合并后的核心路由逻辑：CPU 与 XPU 统一走 AutoAWQMarlinLinearMethod。
class AutoAWQConfig(QuantizeMethodConfig):
    def get_quant_method(self, layer, prefix):
        if isinstance(layer, LinearBase) or (isinstance(layer, ParallelLMHead) and self.lm_head_
            quantized):
            if is_layer_skipped(prefix, self.modules_to_not_convert, self.packed_modules_mapping,
                skip_with_substr=True):
                return UnquantizedLinearMethod()

            # 统一 CPU/XPU：都走 choose_mp_linear_kernel 选择 WNA16 内核。
            # AWQ 权重会在 process_weights_after_loading 时被转换为标准 GPTQ 格式。
            if current_platform.is_cpu() or current_platform.is_xpu():
                return AutoAWQMarlinLinearMethod(self)

            # CUDA 分支：Marlin 支持性检查，失败则回落到基础 AWQ 实现。
            use_marlin = (
                not envs.VLLM_BATCH_INVARIANT
                and current_platform.is_cuda()
                and check_marlin_supported(
                    self.quant_type, self.group_size, self.zero_point
                )
            )
            if use_marlin:
                if not check_marlin_supports_layer(
                    layer, self.group_size, allow_tile_padding=True
                ):
                    return UnquantizedLinearMethod()
```

```

        logger.warning_once(
            "Layer '%s' is not supported by AutoAWQMarlin. "
            "Falling back to unoptimized AWQ kernels.",
            prefix,
        )
        return AutoAWQLinearMethod(self)
    quant_method = AutoAWQMarlinLinearMethod(self)
    quant_method.input_dtype = get_marlin_input_dtype(prefix)
    return quant_method
return AutoAWQLinearMethod(self)
...

```

```

class AutoAWQMarlinLinearMethod(LinearMethodBase):
    def __init__(self, quant_config: AutoAWQConfig) -> None:
        self.quant_config = quant_config
        self.quant_type = scalar_types.uint4
        self.input_dtype = None

    # CPU/XPU 使用专用 WNA16 内核，不经过 Marlin 校验；
    # CUDA 才需要校验 Marlin 支持性。
    if not (current_platform.is_cpu() or current_platform.is_xpu()):
        verify_marlin_supported(
            quant_type=self.quant_config.quant_type,
            group_size=self.quant_config.group_size,
            has_zp=self.quant_config.zero_point,
        )

```

.buildkite/intel_jobs/quantization.yaml

在 Intel XPU 量化 CI 中补充 AWQ 测试，覆盖新的内核路由路径，防止回归。

```

# 在原有 per-token KV cache 测试之后串联 AWQ 测试，
# 确保 XPU 上新的 choose_mp_linear_kernel 路由路径被 CI 覆盖。
- >-
  bash .buildkite/scripts/hardware_ci/run-intel-test.sh
  'VLLM_TEST_FORCE_LOAD_FORMAT=auto pytest -v -s tests/quantization/test_per_token_kv_
  cache.py --deselect="tests/quantization/test_per_token_kv_cache.py::test_triton_unified_
  attention_per_token_head_scale[int4-16-128-num_heads0-seq_lens1]"
  && VLLM_TEST_FORCE_LOAD_FORMAT=auto pytest -v -s tests/quantization/test_auto_awq.py'

```

评论区精华

该 PR 来自 fork，claude[bot] 自动 review 被禁用；合并者 jikunshang 在评论区多次触发 [/ci retry](#) 与 [/ci run](#) 重新执行 Buildkite CI（第 82162、82471、82609 次构建），最终批准合并。没有实质性的技术讨论或设计争议。

- fork 自动 review 禁用与 CI 重试 (other): 没有技术性讨论；jikunshang 最终批准合并。

风险与影响

- 风险:

1. 权重格式转换一致性: XPU 从 AWQUtils.repack + transpose_onednn_woq_format 切换到 AutoAWQMarlinLinearMethod 内的 GPTQ 格式转换, 两者输出必须逐位等价, 否则会产生数值偏差。当前测试覆盖 Qwen2-1.5B AWQ 端到端, 但未覆盖 group_size=-1 等边界配置。
2. Marlin 验证跳过范围扩大: XPU 与 CPU 一样跳过 verify_marlin_supported, 若未来 XPU 内核 (XPUwNa16LinearKernel) 对某些量化参数组合支持不完整, 可能静默产生错误结果。
3. 依赖内核可用性: 删除专有 int4_gemm_w4a16 调用后, XPU AWQ 完全依赖 choose_mp_linear_kernel 能正确选择并加载 XPUwNa16LinearKernel; 若内核缺失或加载失败, 将直接报错而非回退。
4. CI 稳定性: test_auto_awq.py 被串接到现有 Intel quantization job 命令中, 若测试在 XPU 上偶发不稳定, 会拉长 CI 时间并增加失败噪音。- 影响: 用户侧: Intel XPU 上使用 AWQ 量化模型的用户将统一走与 CPU 一致的 kernel 选择路径, 规避 #41469 中 CUDA-only 算子导致的问题, 推理行为与 CPU 对齐。代码侧: 删除约 70 行专用分支, 减少平台特化代码维护成本, 未来 AWQ kernel 选择只需维护 AutoAWQMarlinLinearMethod 一套逻辑。团队 /CI: Intel XPU 量化 CI 从仅覆盖 per-token KV cache 扩展到覆盖 AWQ, 后续 XPU 量化回归能力增强。影响范围集中在量化层和 Intel CI 配置, 不涉及 CUDA/CPU 路径行为变化。- 风险标记: XPU 专属路径移除, 跳过 Marlin 验证, 内核选择依赖 WNA16, 新增 CI 测试覆盖

关联脉络

- PR #42727 AWQ config unification (PR 正文提及, 完整标题未提供): 该 PR 曾统一 AWQ 配置并修复 CPU, 但遗留 XPU 专用方法; 本次变更是其后续, 消除 XPU 平台分歧。