

PR #47521 完整报告

vllm-project/vllm

[Quantization][INC][ARK] Support INT2 XPU WOQ Linear

合并时间: 2026-07-14 08:29

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/47521>

执行摘要

- 一句话: 支持 Intel XPU INT2 权重量化
- 推荐动作: 值得关注的设计决策是 int2 强制依赖 ARK 而不提供 fallback, 这简化了代码但限制了可移植性。评审中的模型规模和版本锁定讨论也具参考价值。建议阅读 `inc_wna16_scheme.py` 的条件逻辑和测试覆盖策略。

功能与动机

在 Intel XPU 上支持更低位宽的模型量化, 利用 ARK 后端加速 int2 线性层计算, 扩展 vLLM 的量化能力以适配硬件特性。PR 标题和 body 中提供了使用 Qwen3-8B-w2g64-for-ut 模型的示例命令, 表明目标是使其能运行 int2 量化模型。

实现拆解

1. 扩展支持位宽集合: 在 `inc_wna16_scheme.py` 中定义 `XPU_WNA16_SUPPORTED_BITS = {2,4}`, 并将条件从 `layer_config.bits == 4` 改为 `layer_config.bits in XPU_WNA16_SUPPORTED_BITS`, 使 int2 配置也能进入分支。
2. 强制 int2 依赖 ARK: 当 `layer_config.bits == 2` 且 ARK 不可用时, 不再像 4-bit 那样 fallback 到 `INCXPULinearMethod`, 而是抛出 `NotImplementedError`, 明确 int2 必须使用 ARK 后端。这体现在 `get_linear_method` 中新增的 `elif` 分支。
3. 测试覆盖: 在 `test_auto_round.py` 中添加 5 个新测试函数:
`test_inc_config_from_config_accepts_xpu_int2` (验证 config parser 正确解析 int2 参数)
`test_wna16_xpu_int2_prefers_ark_when_available` (ARK 可用时返回 `INCARKLinearMethod`)
`test_wna16_xpu_int2_requires_ark_when_unavailable` (ARK 不可用时抛出预期异常)
`test_wna16_xpu_int2_unsupported_config_still_raises` (非对称 int2 仍被拒绝)
`test_inc_ark_linear_method_xpu_int2_create_weights` (验证 `create_weights` 调用顺序)。同时, 修改现有测试 `test_wna16_xpu_unsupported_config_still_raises` 使用 int2 配置以反映新逻辑。还添加了一个端到端模型测试 (`Intel/Qwen3-8B-w2g64-for-ut`), 使用 `tp=2`, 并指定 `block_size`、`gpu_memory_utilization` 等参数。
4. 依赖锁定: 将 `requirements/xpu.txt` 中的 `auto_round_lib` 从 `>=0.14.0` 改为 `==0.14.1`, 确保可重复构建。

关键文件:

- vllm/model_executor/layers/quantization/inc/schemes/inc_wna16_scheme.py (模块 量化方案; 类别 source; 类型 core-logic; 符号 INCWna16Scheme, get_linear_method, XPU_WNA16_SUPPORTED_BITS) : 核心逻辑: 扩展位宽支持并强制 int2 依赖 ARK
- tests/quantization/test_auto_round.py (模块 量化测试; 类别 test; 类型 test-coverage ; 符号 test_inc_config_from_config_accepts_xpu_int2, test_wna16_xpu_int2_prefers_ark_when_available, test_wna16_xpu_int2_requires_ark_when_unavailable, test_wna16_xpu_int2_unsupported_config_still_raises) : 全面测试覆盖: 新增 5 个 int2 单元测试并修改现有测试
- requirements/xpu.txt (模块 依赖管理; 类别 infra; 类型 configuration) : 锁定 auto_round_lib 版本确保兼容性

关键符号: INCWna16Scheme.get_linear_method, test_inc_config_from_config_accepts_xpu_int2, test_wna16_xpu_int2_prefers_ark_when_available, test_wna16_xpu_int2_requires_ark_when_unavailable, test_wna16_xpu_int2_unsupported_config_still_raises, test_inc_ark_linear_method_xpu_int2_create_weights

关键源码片段

vllm/model_executor/layers/quantization/inc/schemes/inc_wna16_scheme.py

核心逻辑: 扩展位宽支持并强制 int2 依赖 ARK

```
# vllm/model_executor/layers/quantization/inc/schemes/inc_wna16_scheme.py
from typing import TYPE_CHECKING
from vllm.logger import init_logger
from vllm.model_executor.layers.quantization.auto_awq import AutoAWQConfig
from vllm.model_executor.layers.quantization.auto_gptq import AutoGPTQConfig
from vllm.platforms import current_platform
from vllm.scalar_type import scalar_types
from ..inc_linear import INCLinearMethod
from .inc_scheme import INCScheme
```

```
if TYPE_CHECKING:
    import torch
    from ..config_parser import INCLayerConfig
    from ..inc import INCConfig
```

```
logger = init_logger(__name__)
```

```
# XPU 上支持的 WNA16 位宽, 现已包含 2-bit
XPU_WNA16_SUPPORTED_BITS = {2, 4}
```

```
class INCWna16Scheme(INCScheme):
    @staticmethod
    def can_handle(layer_config: "INCLayerConfig") -> bool:
```

```

        return layer_config.is_wna16_int

def get_linear_method(
    self,
    config: "INCCConfig",
    layer: "torch.nn.Module",
    prefix: str,
    layer_config: "INCLayerConfig",
):
    del config, layer
    if current_platform.is_xpu():
        # 现在 bits 可以是 2 或 4, 且必须为对称量化
        if layer_config.bits in XPU_WNA16_SUPPORTED_BITS and layer_config.sym:
            from .inc_ark_ops import get_ark_state
            from .inc_wna16_linear import (
                INCARKLinearMethod,
                INCXPULinearMethod,
            )
            is_ark_available, ark_error, _, _ = get_ark_state()
            if is_ark_available:
                return INCLinearMethod(INCARKLinearMethod(layer_config))
            elif layer_config.bits == 2:
                # int2 必须使用 ARK 后端, 无 fallback
                raise NotImplementedError(
                    "INC int2 on XPU requires the ARK backend. "
                    f"Layer: {prefix}. "
                    f"auto_round_kernel unavailable: "
                    f"{ark_error or 'unknown error'}"
                )
            # 4-bit 时 ARK 不可用回退到 XPU 原生实现
            logger.debug(
                "ARK backend is unavailable for layer %s; "
                "falling back to the default XPU INC path. Error: %s",
                prefix,
                ark_error or "unknown error",
            )
            return INCLinearMethod(INCXPULinearMethod(layer_config))
            raise NotImplementedError(f"INC on XPU: unsupported config {layer_config}")

# CPU 分支保持 4-bit 不变
if current_platform.is_cpu() and layer_config.is_gptq:
    if layer_config.bits == 4 and layer_config.sym:
        from .inc_ark_ops import get_ark_state
        from .inc_wna16_linear import (
            INCARKLinearMethod,
            INCWNA16LinearScheme,
        )
        is_ark_available, ark_error, _, _ = get_ark_state()
        if is_ark_available:

```

```

        return INCLinearMethod(INCARKLinearMethod(layer_config))
    logger.debug(
        "ARK backend is unavailable for layer %s; "
        "falling back to the default CPU INC path. Error: %s",
        prefix,
        ark_error or "unknown error",
    )
    return INCLinearMethod(INCWNA16LinearScheme(layer_config))
    raise NotImplementedError(f"INC on CPU: unsupported config {layer_config}")

```

tests/quantization/test_auto_round.py

全面测试覆盖：新增 5 个 int2 单元测试并修改现有测试

tests/quantization/test_auto_round.py (节选)

验证 INCCConfig 能正确解析 int2 原始配置

```

def test_inc_config_from_config_accepts_xpu_int2() -> None:
    # 构建模拟的 int2 原始配置字典
    def _make_int2_raw_config(**overrides) -> dict[str, object]:
        kwargs = {
            "bits": 2,
            "group_size": 64,
            "sym": True,
            "data_type": "int",
            "quant_method": "auto-round",
        }
        kwargs.update(overrides)
        return kwargs

    # 通过 INCCConfig.from_config 解析
    config = INCCConfig.from_config(_make_int2_raw_config())

    # 验证字段正确映射
    assert config.weight_bits == 2
    assert config.group_size == 64
    assert config.sym is True
    assert config.data_type == "int"
    assert config.packing_format == "auto_round:auto_gptq"
    assert config.backend == "auto"

```

验证当 ARK 可用时 int2 选用 ARK 路径

```

def test_wna16_xpu_int2_prefers_ark_when_available(monkeypatch) -> None:
    class DummyQuantLinear:
        pass

    monkeypatch.setattr(current_platform, "is_xpu", lambda: True)
    monkeypatch.setattr(current_platform, "is_cpu", lambda: False)
    monkeypatch.setattr(

```

```

        "vllm.model_executor.layers.quantization.inc.schemes.inc_ark_ops.get_ark_state",
        lambda: (True, None, object(), DummyQuantLinear),
    )

    method = INCWna16Scheme().get_linear_method(
        make_config(weight_bits=2, group_size=64),
        object(),
        "layer",
        make_layer_config(bits=2, group_size=64),
    )

    assert isinstance(method, INCLinearMethod)
    assert isinstance(method.scheme, INCARKLinearMethod)

```

评论区精华

评审过程中有多条讨论：

- yiliu30建议定义常量替代硬编码集合（已采纳）；要求添加模型级测试（已补充）；指出原 72B 模型过大，建议换小模型（已改为 8B）。
- yma11质疑 int2 不可用时报错而非 fallback 的设计，Zhenzhong1解释 int2 必须依赖 ARK，vLLM XPU kernel 不支持 int2，故意报错以明确要求，其评论得到认可。
- jikunshang建议将 if 改为 elif 提高可读性（已修改）。
- yma11询问 4-bit 下 ARK 与 vLLM kernel 性能对比，开发者承诺后续提供 benchmark。
- 定义支持位宽常量 (design): 接受建议，添加 `XPU_WNA16_SUPPORTED_BITS` 常量。
- 添加模型级集成测试 (testing): 开发者添加了 Intel/Qwen3-8B-w2g64-for-ut 模型测试，并使用 skipif 避免在多卡不足时运行。
- int2 无 ARK 时的处理策略 (design): 设计被接受，保持 int2 强制依赖 ARK 的语义。
- if 改 elif 提高可读性 (style): 接受建议，将 `if layer_config.bits == 2` 改为 elif。

风险与影响

- 风险：主要风险是 int2 强制依赖 ARK，若用户在无 ARK 环境下加载 int2 模型将直接失败，但这是设计意图，且提供了清晰错误信息。此外，依赖版本锁定可能影响其他使用 `auto_round_lib` 的组件，但锁定版本为 bugfix 级别，兼容性风险低。新增模型测试（8B，tp=2）虽已缩减规模，仍可能占用 CI 资源，但带有 skipif 条件仅在多 XPU 上运行，影响可控。
- 影响：影响范围局限于 Intel XPU 用户，特别是需要 int2 量化的场景。现有 4-bit 量化路径不受影响（测试已覆盖）。用户必须在环境中正确配置 ARK 才能使用 int2 功能。团队内部需要维护额外的测试用例和依赖版本。由于改动集中在单一 scheme 文件且测试充分，回归风险较低。
- 风险标记：int2 强制依赖 ARK，新增模型测试资源占用，依赖版本锁定

关联脉络

- 暂无明显关联 PR