

PR #40601 完整报告

vllm-project/vllm

[quant][autoround]Refactor INC quantization into package with INCScheme orchestrator

合并时间: 2026-06-17 21:51

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/40601>

执行摘要

- 一句话: 重构 INC 量化模块为 scheme-based 包结构
- 推荐动作: 值得精读。重点关注 scheme 接口设计如何支持未来新量化类型、config_parser 的 regex 匹配实现、ARK 后端的可用性检测与回退策略、以及平台差异下 MoE 方法的处理。

功能与动机

按照 compressed_tensors 的风格, 将 monolithic inc.py 替换为 scheme-based 分发的 inc/ 包, 以便在近期内添加更多量化方案。见 Issue #37979。

实现拆解

1. 创建包架构: 将原 inc.py 拆分为 inc/ 包, 包含 init.py、inc.py、config_parser.py、inc_linear.py 及 schemes/ 子包, 保留 INCConfig 作为配置入口, 将层配置解析委托给 INCConfigParser。
2. 定义 scheme 抽象接口: 在 inc/schemes/inc_scheme.py 中定义 INCScheme (含 can_handle、get_linear_method、get_moe_method、get_kvcache_method) 和 INCLinearScheme (含 create_weights、apply_weights) 抽象基类, 每种量化类型一个子类。
3. 实现配置解析和首个 scheme: INCConfigParser 支持 exact match 和 regex match 的层配置解析, 返回 INCLayerConfig frozen dataclass; INCWna16Scheme 作为首个 scheme, 在 get_linear_method 中根据平台 (XPU/CPU) 和 ARK 可用性选择 ARK 或原生方法, get_moe_method 中对 XPU/CPU 返回 UnquantizedFusedMoEMethod, 其他平台使用 GPTQ 或 AWQ Marlin。
4. 清理旧文件并更新测试: 删除旧 inc.py (-794 行), 更新各 init.py 导出, 重写 test_auto_round.py (+750 行) 增加新单元测试和模型测试, 在 .buildkite/intel_jobs/test-intel.yaml 中添加 XPU 量化测试步骤。

关键文件:

- vllm/model_executor/layers/quantization/inc.py (模块 INC 量化; 类别 source; 类型 deletion; 符号 INCConfig, init, get_name, from_config): 旧入口文件, 全部删除 (-794 行), 改为新包结构, 是整个重构的靶点。

- `vllm/model_executor/layers/quantization/inc/inc.py` (模块 INC 量化; 类别 `source`; 类型 `data-contract`; 符号 `INCCConfig`, `init`, `get_quant_method`, `get_layer_config`) : 新的 INC 入口模块, 封装 `INCCConfig` 及 `get_quant_method`、`get_layer_config`、`apply_vllm_mapper`, 通过 `config_parser` 委托层配置解析。
- `vllm/model_executor/layers/quantization/inc/config_parser.py` (模块 配置解析; 类别 `source`; 类型 `data-contract`; 符号 `INCLayerConfig`, `INCCConfigParser`, `resolve`, `_resolve_raw`) : 新增配置解析模块, 定义 `INCLayerConfig` frozen dataclass (含 `is_gptq/is_awq/is_wna16_int` 等属性) 和 `INCCConfigParser.resolve` 方法, 支持逐层 `exact/regex` 匹配。
- `vllm/model_executor/layers/quantization/inc/schemes/inc_scheme.py` (模块 Scheme 接口; 类别 `source`; 类型 `data-contract`; 符号 `INCScheme`, `INCLinearScheme`, `can_handle`, `get_linear_method`) : 定义 INC 量化方案的抽象基类 `INCScheme` (`can_handle`、`get_linear_method`、`get_moe_method`、`get_kvcache_method`) 和 `INCLinearScheme` (`create_weights`、`apply_weights`) , 所有具体方案必须实现。
- `vllm/model_executor/layers/quantization/inc/schemes/inc_wna16_linear.py` (模块 WNA16 线性; 类别 `source`; 类型 `data-contract`; 符号 `INCWNA16LinearScheme`, `get_ark_state`, `_build_inner_method`) : `INCWNA16LinearScheme` 实现了 `INCLinearScheme`, 根据不同后端 (GPTQ/AWQ/Marlin/ARK) 构建内部线性方法, 包含 `get_ark_state` 检测 ARK 可用性。

关键符号: `INCCConfig.get_quant_method`, `INCCConfigParser.resolve`, `INCCConfigParser._resolve_raw`, `INCScheme.can_handle`, `INCScheme.get_linear_method`, `INCScheme.get_moe_method`, `INCWna16Scheme.get_linear_method`, `INCWna16Scheme.get_moe_method`, `INCWNA16LinearScheme._build_inner_method`, `get_ark_state`, `resolve_scheme`

关键源码片段

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

新的 INC 入口模块, 封装 `INCCConfig` 及 `get_quant_method`、`get_layer_config`、`apply_vllm_mapper`, 通过 `config_parser` 委托层配置解析。

```
# vllm/model_executor/layers/quantization/inc/inc.py (new) # 包入口, 继承
QuantizationConfig, 通过 config_parser 将层配置解析委托给 INCCConfigParser
from fractions import Fraction
from typing import TYPE_CHECKING, Any
import torch
from .config_parser import INCCConfigParser
from .schemes.factory import resolve_scheme

class INCCConfig(QuantizationConfig):
    """Config class for Intel Neural Compressor (INC)."""
    def __init__(self, ...):
        # ... 参数校验, 赋值
        self.config_parser = INCCConfigParser(self) # 将解析逻辑分离给专用 parser

    def get_quant_method(self, layer, prefix):
        """主分派入口: 先检查 layer 类型, 再通过 scheme 工厂获得方法"""
        from vllm.model_executor.layers.fused_moe import FusedMoE, RoutedExperts
        from vllm.model_executor.layers.linear import LinearBase
        from vllm.model_executor.layers.vocab_parallel_embedding import ParallelLMHead

        # 1. 解析层配置为 INCLayerConfig
        layer_config =
```

```

self.config_parser.resolve(layer, prefix)      # 2. 未量化层早期返回      if not
layer_config.quantized:      if isinstance(layer, (LinearBase, ParallelLMHead)):
    from vllm.model_executor.layers.linear import UnquantizedLinearMethod
    return UnquantizedLinearMethod()          if isinstance(layer, (FusedMoE,
RoutedExperts)):      from vllm.model_executor.layers.fused_moe import
UnquantizedFusedMoEMethod          return
UnquantizedFusedMoEMethod(layer.moe_config)    # 3. 通过 scheme 工厂选择对应量
化方案的 Linear / MoE 方法      scheme = resolve_scheme(layer_config)      if
isinstance(layer, (LinearBase, ParallelLMHead)):      return
scheme.get_linear_method(self, layer, prefix, layer_config)      if isinstance(layer,
(FusedMoE, RoutedExperts)):      return scheme.get_moe_method(self, layer,
prefix, layer_config)      raise NotImplementedError(...) (注：实际代码行数较多，此处
展示核心流程，注释已内联。)

```

vllm/model_executor/layers/quantization/inc/config_parser.py

新增配置解析模块，定义 INCLayerConfig frozen dataclass (含 is_gptq/is_awq/is_wna16_int 等属性) 和 INCConfigParser.resolve 方法，支持逐层 exact/regex 匹配。

```

# vllm/model_executor/layers/quantization/inc/config_parser.py (new) # 定义层配置的数据
结构 (INCLayerConfig) 和解析器 (INCConfigParser) @dataclass(frozen=True)
class INCLayerConfig: bits: int group_size: int sym: bool packing_format:
str backend: str data_type: str quantized: bool # 通过 packing_format 或
backend 判断量化子类型 @property def is_gptq(self) -> bool: return "gptq" in
self.packing_format or "gptq" in self.backend @property def is_awq(self) -> bool:
return "awq" in self.packing_format or "awq" in self.backend @property
def is_wna16_int(self) -> bool: return self.data_type == "int" and self.quantized
class INCConfigParser: def __init__(self, config: "INCConfig"): self._config =
config def resolve(self, layer, layer_name) -> INCLayerConfig: bits, group_size,
sym = self._resolve_raw(layer, layer_name) return INCLayerConfig(
bits=bits, group_size=group_size, sym=sym,
packing_format=self._config.packing_format, backend=self._config.backend,
data_type=self._config.data_type, quantized=bits < 16)
def _resolve_raw(self, layer, layer_name) -> tuple[int, int, bool]: # 支持
extra_config 中的 exact key 和 regex 匹配，未匹配则返回全局默认值
def get_config(name, quantized=True): if not self._config.extra_config:
return (self._config.weight_bits if quantized else 16,
self._config.group_size if quantized else -1, self._config.sym if
quantized else True) if name in self._config.extra_config: cfg =
self._config.extra_config[name] return (cfg.get("bits",
self._config.weight_bits if quantized else 16),
self._config.group_size if quantized else -1),
self._config.group_size if quantized else -1),
self._config.sym if quantized else True)) for pattern, cfg in
self._config.extra_config.items(): if not isinstance(pattern, str) or not any(c

```

```

in set(r"*+?^$()\[\]{}|\\") for c in pattern):
    if re.search(re.compile(pattern), name):
        ..), ...)
    except re.error:
        ... (注释内联说明解析回退逻辑。)

```

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

定义 INC 量化方案的抽象基类 INCScheme (can_handle、get_linear_method、get_moe_method、get_kvcache_method) 和 INCLinearScheme (create_weights、apply_weights)，所有具体方案必须实现。

```

# vllm/model_executor/layers/quantization/inc/schemes/inc_scheme.py (new) # 定义两个
抽象基类: INCScheme (高层方案调度) 和 INCLinearScheme (线性层方法)
class INCScheme(ABC): """每个量化类型一个子类, 通过 can_handle 判断是否适合当前层配
置。"""
    @staticmethod
    @abstractmethod
    def can_handle(layer_config: "INCLayerConfig") -> bool:
        raise NotImplementedError
    @abstractmethod
    def get_linear_method(self, config, layer, prefix, layer_config) -> "LinearMethodBase": """
必须实现: 返回线性层的量化方法。"""
        raise NotImplementedError
    def get_moe_method(self, config, layer, prefix, layer_config) -> "FusedMoEMethodBase |
None": """可选实现 MoE 方法, 默认抛 NotImplementedError。"""
        raise
        NotImplementedError(f"{type(self).__name__} does not support MoE layers.")
    def get_kvcache_method(self, config, layer, prefix, layer_config) -> "
QuantizationMethods": """可选实现 KV cache 量化, 默认抛 NotImplementedError。"""
        raise NotImplementedError(...)
class INCLinearScheme(ABC): """线性层内部方法的抽象
, 包含权重创建、加载和应用。"""
    @classmethod
    @abstractmethod
    def get_min_capability(cls) -> int: ...
    @abstractmethod
    def create_weights(self,
layer, ...): ...
    @abstractmethod
    def process_weights_after_loading(self, layer): ...
    @abstractmethod
    def apply_weights(self, layer, x, bias): ... (方案工厂
resolve_scheme 遍历所有 INCScheme 子类, 调用 can_handle 选择首个匹配。)

```

评论区精华

- MoE 方法兼容性: gemini-code-assist 指出返回 UnquantizedLinearMethod 给 FusedMoE 会导致类型签名不匹配, 作者修复为返回 UnquantizedFusedMoEMethod(layer, moe_config)。
- 命名规范对齐: jikunshang 要求遵循 compressed_tensor 风格, 作者融合了 linear 和 moe 方法于 inc_wna16_linear。
- 类名调整: jikunshang 建议 renaming schemes to schema, 作者采纳。
- 配置解析器命名: hshen14 建议将 resolver 重命名为 parser, 作者改为 config_parser.py。
- Future attention 量化: hshen14 询问 attention 支持, 作者表示后续通过 AutoroundKVCacheMethod 处理, 当前不涉及。
 - MoE 层返回错误方法类型 (correctness): 作者修复为返回 UnquantizedFusedMoEMethod(layer, moe_config)。

- 命名规范：遵循 compressed_tensor 风格 (design): 作者调整文件命名，融合 linear 和 moe 方法于 inc_wna16_linear。
- 类名：scheme vs schema (style): 作者采纳并重命名。
- 配置解析器命名：resolver vs parser (style): 作者采纳改为 config_parser.py。
- Attention 量化支持计划 (design): 作者表示后续通过 AutoroundKVCacheMethod 支持，当前不涉及。

风险与影响

- 风险：
 1. 兼容性风险：旧 inc.py 中某些配置键（如 block_name_to_quantize 拆分路径）在新 config_parser 中行为可能变化，需验证已有模型配置。
 2. 依赖风险：ARK 后端自动选择需要 auto-round-lib >= 0.13.3，否则启动时崩溃（用户 urakozz 报告过），代码中未显式检查版本。
 3. MoE 返回值风险：最初返回错误方法类型（已修复），但仍需确认在所有平台组合下的正确性。
 4. 回归风险：重构核心量化分派逻辑，可能影响所有使用 INC 量化的模型（如 Qwen3、LLaMa 等）的推理结果。 - 影响：用户：对外部透明，但需要升级 auto-round-lib 依赖 (>= 0.13.3)。系统：架构更易扩展新量化方案，长期降低维护成本。团队：INC 模块与 compressed_tensors 架构对齐，降低跨模块理解成本。 - 风险标记：缺少 auto-round-lib 版本检查，依赖外部 kernel 库，核心重构风险

关联脉络

- PR #37979 [RFC]: Intel Quantization Support Roadmap (H1 2026): 本 PR 是该 RFC 中 'Architectural Cleanup' 项的具体实现，用于解耦 INC 量化模块。