

PR #47091 完整报告

vllm-project/vllm

[Bugfix] [Gemma4] Fix Gemma4 MTP draft model layers ignoring quant_config

合并时间: 2026-07-06 21:04

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/47091>

执行摘要

- 一句话: 修复 Gemma4 MTP 草稿模型量化配置未传播问题
- 推荐动作: PR 值得合并, 修复了一个关键 Bug, 且改动量小、有测试验证。建议查看 `get_draft_quant_config` 的实现以确保其正确性。

功能与动机

量化后的 Gemma4 MTP 草稿检查点草稿 Token 接受率为 0%, 因为所有层都硬编码了 `quant_config=None`, 导致量化权重未被正确加载。修复后接受率从 0% 提升至 65.9%, 平均接受长度从 1.00 增至 3.64。

实现拆解

1. 导入 `get_draft_quant_config`: 在 `gemma4_mtp.py` 的 import 中从 `vllm.model_executor.models.utils` 添加该函数, 用于获取草稿模型专属的量化配置。
2. 修复 `Gemma4MTPAttention`: 将 `q_proj` 和 `o_proj` 的 `quant_config` 从 `None` 改为构造参数传入的 `quant_config`。
3. 修复 `Gemma4MTPDecoderLayer`: 将 MLP 构造时的 `quant_config` 从 `None` 改为传入的 `quant_config`。
4. 修复 `Gemma4MultiTokenPredictor`: 使用 `get_draft_quant_config(vllm_config)` 获取草稿模型量化配置, 并传给 `embed_tokens`、`pre_projection`、`post_projection` 及所有 `Gemma4MTPDecoderLayer`; 此前错误使用了目标模型的 `vllm_config.quant_config`。
5. 修复 `Gemma4MTP`: 同样使用 `get_draft_quant_config` 获取草稿量化配置, 并传给 `lm_head` 的 `ParallelLMHead`。

关键文件:

- `vllm/model_executor/models/gemma4_mtp.py` (模块 模型层; 类别 `source`; 类型 `data-contract`; 符号 `Gemma4MTPAttention`, `Gemma4MTPDecoderLayer`, `Gemma4MultiTokenPredictor`, `Gemma4MTP`): 核心文件, 修复了 MTP 草稿模型所有层量化配置硬编码问题

关键符号: `Gemma4MTPAttention.init`, `Gemma4MTPDecoderLayer.init`, `Gemma4MultiTokenPredictor.init`, `Gemma4MTP.init`

关键源码片段

vllm/model_executor/models/gemma4_mtp.py

核心文件，修复了 MTP 草稿模型所有层量化配置硬编码问题

gemma4_mtp.py 中关键变更：修复量化配置传播

```
from .utils import (
    AutoWeightsLoader,
    WeightsMapper,
    extract_layer_index,
    get_draft_quant_config, # 新增导入：用于获取草稿模型专属量化配置
    maybe_prefix,
)

class Gemma4MultiTokenPredictor(nn.Module):
    def __init__(self, *, vllm_config: VllmConfig, prefix: str = ''):
        config = vllm_config.speculative_config.draft_model_config.hf_config
        text_config = _get_text_config(config)
        quant_config = get_draft_quant_config(vllm_config) # 使用草稿模型配置，而非目标模型配置

        ...
        self.embed_tokens = VocabParallelEmbedding(
            self.vocab_size,
            self.hidden_size,
            quant_config=quant_config, # 之前为 None
            prefix=f'{prefix}.embed_tokens',
        )
        self.pre_projection = ColumnParallelLinear(
            2 * self.backbone_hidden_size,
            self.hidden_size,
            bias=False,
            gather_output=True,
            quant_config=quant_config, # 之前为 None
            ...
        )
        self.post_projection = RowParallelLinear(
            self.backbone_hidden_size,
            self.hidden_size,
            bias=False,
            input_is_parallel=False,
            quant_config=quant_config, # 之前为 None
            ...
        )
        self.layers = nn.ModuleList(
            Gemma4MTPDecoderLayer(
                text_config,
                cache_config=vllm_config.cache_config,
                quant_config=quant_config, # 之前误用 vllm_config.quant_config
                ...
            )
        )
```

```

        for _ in range(self.num_mtp_layers)
    )

class Gemma4MTPAttention(nn.Module):
    def __init__(self, ..., quant_config: QuantizationConfig | None = None, ...):
        ...
        self.q_proj = ColumnParallelLinear(
            hidden_size,
            self.total_num_heads * self.head_dim,
            bias=config.attention_bias,
            quant_config=quant_config, # 之前为 None
            ...
        )
        self.o_proj = RowParallelLinear(
            self.total_num_heads * self.head_dim,
            hidden_size,
            bias=config.attention_bias,
            quant_config=quant_config, # 之前为 None
            ...
        )

class Gemma4MTPDecoderLayer(nn.Module):
    def __init__(self, config, ..., quant_config: QuantizationConfig | None = None, ...):
        ...
        self.mlp = Gemma4MLP(
            hidden_size=self.hidden_size,
            intermediate_size=text_config.intermediate_size,
            hidden_activation=text_config.hidden_activation,
            quant_config=quant_config, # 之前为 None
            ...
        )

```

评论区精华

- 暂无高价值评论线程

风险与影响

- 风险：变更仅局限于 Gemma4 MTP 模型模块（gemma4_mtp.py），影响范围有限。非量化模型行为不变（get_draft_quant_config 正常返回 None 时回退）。主要风险是如果 get_draft_quant_config 实现有 bug 或返回错误配置，可能导致量化加载异常，但已有测试覆盖。
- 影响：对用户：修复了 Gemma4 量化草稿模型在推测解码中完全失效的 bug，恢复正确行为。对系统：无性能影响，仅修改量化配置传播逻辑。对团队：新增了一个关键 Bugfix，量化模型用户受益。
- 风险标记：核心路径变更

关联脉络

- 暂无明显关联 PR