

PR #52134 完整报告

vllm-project/vllm

[Docs] Fix `WhisperEncoderLayer.forward` docstring in `dots3\_note`

合并时间: 2026-08-13 17:32

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/52134>

执行摘要

- 一句话: 修复音频编码器 docstring, 消除 griffe 文档警告
- 推荐动作: 值得浏览但不必精读: 可作为『docstring 与签名对齐』的维护样例, 快速理解如何排查 griffe 文档警告并修正继承自上游的过时 docstring。注意其本质是数据契约 (返回类型注解) 的文档化修正, 而非契约本身的运行时变更。

功能与动机

docs 构建对 vllm/models/dots3_note/nvidia/audio_encoder.py 报出两条 griffe 警告 (Parameter 'attention_mask' / 'layer_head_mask' does not appear in the function signature)。docstring 从上游 HF Whisper 继承后未适配本层实际签名——该层接收 packed 变长输入 (cu_seqlens_、max_seqlen_) 与旋转位置编码 (rotary_cos/rotary_sin), 而非 attention_mask/layer_head_mask; 返回注解声称 torch.Tensor, 实际却返回 hidden states 与可选 attention weights 组成的元组。PR body 同时说明这不是重复工作: #51342 仅修复 vllm/benchmarks/throughput.py 中的无关注解。

实现拆解

实现按以下 4 步展开:

1. 定位问题来源: 在 vllm/models/dots3_note/nvidia/audio_encoder.py 中定位 WhisperEncoderLayer.forward, 确认 docstring 保留了上游 HF Whisper 的 attention_mask、layer_head_mask 参数说明, 而签名已改为 cu_seqlens_q/cu_seqlens_kv/max_seqlen_q/max_seqlen_kv/output_attentions/rotary_cos/rotary_sin, 这是 griffe 警告的直接原因。
2. 重写 Args 段: 将参数说明逐一对齐到真实签名, 补上 max_seqlen_q、max_seqlen_kv 等此前缺失的说明; 对 hidden_states 同时说明 padded batch ((batch, seq_len, embed_dim)) 与 packed 变长 ((total_tokens, embed_dim)) 两种输入形态。
3. 修正返回契约: 返回注解由 torch.Tensor 改为 tuple[Any, ...], 并新增 Returns 段说明返回 hidden states 以及 output_attentions=True 时附加的 attention weights, 与 forward_flash_attn 的实际返回一致。
4. 验证配套: 无测试、配置或部署配套改动 (运行时零变化); 作者用 griffe 的 Google-style docstring 解析器验证警告消除, pre-commit (mypy、ruff) 通过。

关键文件:

- vllm/models/dots3_note/nvidia/audio_encoder.py (模块 音频编码器; 类别 source; 类型 data-contract; 符号 WhisperEncoderLayer.forward) : 本次唯一变更文件, 重写 WhisperEncoderLayer.forward 的 docstring 并修正返回注解, 消除 griffe 文档构建警告。

关键符号: WhisperEncoderLayer.forward

关键源码片段

vllm/models/dots3_note/nvidia/audio_encoder.py

本次唯一变更文件, 重写 WhisperEncoderLayer.forward 的 docstring 并修正返回注解, 消除 griffe 文档构建警告。

```
def forward(
    self,
    hidden_states: torch.Tensor,
    cu_seqlens_q: torch.Tensor = None,
    cu_seqlens_kv: torch.Tensor = None,
    max_seqlen_q: int | None = None,
    max_seqlen_kv: int | None = None,
    output_attentions: bool = False,
    rotary_cos: torch.Tensor | None = None,
    rotary_sin: torch.Tensor | None = None,
) -> tuple[Any, ...]:
    """
    Args:
        hidden_states: Input to the layer of shape `(batch, seq_len,
            embed_dim)`, or `(total_tokens, embed_dim)` when
            `cu_seqlens_q` is given.
        cu_seqlens_q: Cumulative query sequence lengths for packed
            variable-length input. If `None`, the input is treated as a
            padded batch.
        cu_seqlens_kv: Cumulative key/value sequence lengths for packed
            variable-length input.
        max_seqlen_q: Longest query sequence in the packed batch.
        max_seqlen_kv: Longest key/value sequence in the packed batch.
        output_attentions: Whether to also return the attention weights.
        rotary_cos: Cosine component of the rotary position embedding.
        rotary_sin: Sine component of the rotary position embedding.

    Returns:
        A tuple of the output hidden states, followed by the attention
        weights if `output_attentions` is `True`.
    """
    # 本次文档修正的要点 (无运行时行为变化) :
    # 1. 原 docstring 继承自上游 HF Whisper, 包含 attention_mask 与
    # layer_head_mask 两个本签名不存在的参数, 导致 griffe 警告;
    # 2. 返回注解由 torch.Tensor 更正为 tuple[Any, ...], 因为下方
    # forward_flash_attn 返回 (hidden_states, attn_weights) 元组,
    # 且 output_attentions=True 时元组会包含 attention weights.
```

```
residual = hidden_states
hidden_states = self.self_attn_layer_norm(hidden_states)

hidden_states, attn_weights = self.self_attn.forward_flash_attn(
    hidden_states=hidden_states,
    cu_seqlens_q=cu_seqlens_q,
    cu_seqlens_kv=cu_seqlens_kv,
    max_seqlen_q=max_seqlen_q,
    max_seqlen_kv=max_seqlen_kv,
    output_attentions=output_attentions,
    rotary_cos=rotary_cos,
    rotary_sin=rotary_sin,
)
```

评论区精华

该 PR 没有实质性技术讨论。claude[bot] 自动提示：来自 fork 的 PR 不启用自动 review，维护者可评论 @claude review 触发一次性审查；DarkLight1337 直接批准（无评论）。核心设计决策由作者在 PR body 中说明：照实记录真实签名、删除无效参数说明、返回注解与实现对齐，并明确排除了与 #51342 的重复性。

- fork PR 自动 review 状态 (other): 维护者 DarkLight1337 直接批准，未触发额外审查。

风险与影响

- 风险：变更仅涉及 docstring 与类型注解，不执行任何运行时逻辑，回归风险极低。两个值得留意的点：1) tuple[Any, ...] 中的 Any 依赖文件既有 typing 导入，作者已用 mypy 验证通过；2) docstring 与签名未来仍可能再次漂移——例如 forward 签名再次调整时若未同步文档，griffe 警告会重新出现，本次只是消除当前警告而非建立机制性预防。
- 影响：对用户：无可观察的运行时效影响，但 API 文档、IDE 提示与自动生成文档的准确性得到改善。对系统：docs 构建不再报 griffe 警告，减轻 CI 噪音。对团队：属于低成本文档维护，为 dots3_note 音频模型对外 API 文档的正确性打基础；同时为其他从上游继承 docstring 的模型层提供了修正范例。
- 风险标记：纯文档变更，文档签名漂移隐患

关联脉络

- 暂无明显关联 PR