

PR #48134 完整报告

vllm-project/vllm

[Bugfix][Rust Frontend] Limit chat top_logprobs in responses

合并时间: 2026-07-16 17:41

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/48134>

执行摘要

此 PR 修复了 Rust 前端 chat completion 响应中 `top_logprobs` 未按请求参数截断的 bug。通过新增 `chat_top_logprob_entries` 函数并沿请求链路传递 `top_logprobs` 值, Rust 前端现在与 Python 前端的行为一致: 默认返回空列表、正数 `k` 返回前 `k` 个、`-1` 返回全部。改动集中在 `logprob` 转换核心函数和路由层, 并附带单元测试与集成测试验证。

功能与动机

Rust 前端在 `logprobs=true` 时, 会返回引擎解码的每一个候选 token, 未根据请求参数 `top_logprobs` 截断。这与 Python 前端的响应不一致。PR body 明确列出 Python 的三种行为:

- `top_logprobs` 省略或为 0: 返回空列表 []
- `top_logprobs` 为正数 `k`: 返回前 `k` 个条目
- `top_logprobs` 为 `-1`: 返回所有条目

此 PR 的目标是让 Rust 前端的 chat logprobs 响应完全对齐 Python 服务。

实现拆解

1. 数据结构扩展 (`convert.rs`) 在 `ResponseOptions` 结构体中新增 `output_top_logprobs: i32` 字段, 并在 `prepare_chat_request` 中将请求的 `top_logprobs` 赋值给该字段, 使参数能够沿请求处理链路传递到下游。
2. 截断逻辑实现 (`logprobs.rs`) 新增 `chat_top_logprob_entries` 函数, 根据 `top_logprobs` 值决定截断数: - `-1`: 返回所有条目 - 其他值: 通过 `usize::try_from(top_logprobs).unwrap_or(0)` 转换为截断长度, 并调用 `position.entries.iter().take(limit)` 实现截断

修改 `decoded_logprobs_to_openai_chat` 和 `position_to_chat_logprobs_content` 函数, 新增 `top_logprobs: i32` 参数, 替换原先对 `position.entries.iter()` 的硬编码全量遍历。

1. 路由层参数传递 (`chat_completions.rs`) 在 `collect_chat_completion` 和 `chat_completion_chunk_stream` 中解构 `output_top_logprobs` 并传递给 `decoded_logprobs_to_openai_chat`, 确保非流式和流式响应使用相同的截断逻辑。
2. 测试覆盖 - 单元测试 (`logprobs.rs`): 在 `#[cfg(test)]` 模块中使用 `sample_logprobs` 构造包含 3 个候选条目的位置, 通过 `chat_top_logprobs_len` 辅助函数验证 `top_logprobs=0`、`1`、`-1` 时的截断结果。 - 集成测试 (`tests.rs`): 在现有 `non_stream_chat_includes_logprobs_and_prompt_logprobs` 测试中添加断言, 验证 `top_logprobs` 初始为空列表。

rust/src/server/src/routes/openai/utils/logprobs.rs

核心实现文件：新增 `chat_top_logprob_entries` 函数进行截断逻辑，修改 `decoded_logprobs_to_openai_chat` 和 `position_to_chat_logprobs_content` 函数签名以接收 `top_logprobs` 参数，并附带单元测试 `chat_logprobs_respects_requested_top_logprobs_count`。

```
/// 根据请求的 top_logprobs 值截断候选条目，返回迭代器。
/// - top_logprobs == -1: 返回所有条目
/// - top_logprobs >= 0: 返回前 top_logprobs 个 (0 返回空)
fn chat_top_logprob_entries(
    position: &DecodedPositionLogprobs,
    top_logprobs: i32,
) -> impl Iterator<Item = &DecodedTokenLogprob> {
    let limit = if top_logprobs == -1 {
        position.entries.len()
    } else {
        // 负数会 unwrap_or 为 0，静默处理
        usize::try_from(top_logprobs).unwrap_or(0)
    };
    position.entries.iter().take(limit)
}

/// 将解码后的 logprobs 转换为 OpenAI chat 格式，传入 top_logprobs 控制截断。
pub fn decoded_logprobs_to_openai_chat(
    logprobs: &DecodedLogprobs,
    top_logprobs: i32,
    return_tokens_as_token_ids: bool,
) -> Result<ChatLogProbs, ApiError> {
    let content = logprobs
        .positions
        .iter()
        .map(|pos| position_to_chat_logprobs_content(pos, top_logprobs, return_tokens_as_token_ids))
        .try_collect()?;

    Ok(ChatLogProbs {
        content: Some(content),
    })
}

#[cfg(test)]
mod tests {
    use super::*;

    fn sample_logprobs() -> DecodedLogprobs {
        DecodedLogprobs {
            positions: vec![DecodedPositionLogprobs {
                entries: vec![
```

```

        DecodedTokenLogprob { token_id: 1, token: "A".to_string(), logprob: -0.1, rank: 1
    },
        DecodedTokenLogprob { token_id: 2, token: "B".to_string(), logprob: -1.0, rank: 2
    },
        DecodedTokenLogprob { token_id: 3, token: "C".to_string(), logprob: -2.0, rank: 3
    },
    ],
    },
}

fn chat_top_logprobs_len(top_logprobs: i32) -> usize {
    let chat_logprobs =
        decoded_logprobs_to_openai_chat(&sample_logprobs(), top_logprobs, false)
        .expect("chat logprobs");
    chat_logprobs.content.expect("content")[0].top_logprobs.len()
}

#[test]
fn chat_logprobs_respects_requested_top_logprobs_count() {
    // top_logprobs=0 应返回空列表
    assert_eq!(chat_top_logprobs_len(0), 0);
    // top_logprobs=1 应返回 1 个条目
    assert_eq!(chat_top_logprobs_len(1), 1);
    // top_logprobs=-1 应返回所有条目
    assert_eq!(chat_top_logprobs_len(-1), 3);
}
}

```

评论区精华

该 PR 无人工 review 评论，仅由 BugenZhao 两次批准并合并（“Thanks!”），未提出设计或代码层面的讨论。自动审查机器人 Claude 因 PR 来自 fork 而跳过了审查。

风险与影响

- 向后兼容性：此前 Rust 前端返回所有候选条目，此 PR 后默认（top_logprobs 未指定时）行为取决于请求解析默认值。Python 前端中 top_logprobs 默认为 0，导致返回空列表，这可能会影响依赖全量 top_logprobs 的客户端。但此 PR 正是为了对齐 Python 行为，属于预期修复。
- 负数处理：当 top_logprobs 为负数且非 -1 时，chat_top_logprob_entries 中 `usize::try_from(top_logprobs).unwrap_or(0)` 会静默截断为 0，与 Python 的 `ValueError` 行为不同。虽然 API 文档通常限制为 -1/0/ 正数，但未来可考虑更严格的参数校验。
- 影响范围：仅影响 Rust 前端的 chat completion 响应路径，不影响 Python 前端或核心引擎，无性能或架构风险。

关联脉络

此 PR 是 Rust 前端持续行为对齐的一部分。近期历史 PR 中，[#48738](#) 修复了 mock engine 测试的竞态条件，[#46647](#) 将迭代日志移到前端。这些 PR 共同表明团队正在积极提升 Rust 前端的稳定性与一致性。