

PR #46344 完整报告

vllm-project/vllm

[Frontend] Fix Kimi K2 tool call IDs for required tool choice

合并时间: 2026-06-25 03:59

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/46344>

执行摘要

- 一句话: 修复 Kimi K2 required tool choice ID 格式
- 推荐动作: 值得精读。该 PR 展示了一次典型的 "逻辑下沉 + 清理" 重构: 将特定模型的 ID 格式策略从高层服务移到低层 parser, 使代码职责更单一、便于测试。新增的 count_history_tool_calls 系列函数设计通用, 可用于未来其他模型。对于关注工具调用流程或 parser 架构的开发者来说, 这是一个很好的设计决策参考。

功能与动机

Kimi K2 模型在 `tool_choice="required"` 场景下, tool call ID 未能正确按照 `functions.<function_name>:<idx>` 格式生成, 且未考虑历史 tool call 的计数偏移。PR 修复此问题, 同时将 ID 生成逻辑下沉到 parser 中, 使架构更清晰。

实现拆解

1. 新增工具模块 `vllm/parser/utils.py`: 提供了 `count_tool_calls`、`count_chat_history_tool_calls`、`count_response_history_tool_calls` 和 `count_history_tool_calls` 四个函数, 用于统计历史对话中已存在的 tool call 数量, 支持 `ChatCompletionRequest` 和 `ResponsesRequest` 两种请求类型。
2. 重构 `vllm/parser/abstract_parser.py`: 在 `StreamState` 中新增 `history_tool_call_cnt_initialized` 标志, 避免重复初始化; 新增 `_initialize_history_tool_call_cnt()` 方法 (当 `tool_call_id_type == "kimi_k2"` 时调用 `count_history_tool_calls`) 和 `_make_tool_call_id()` 方法 (生成带递增序号的 ID); 在 `Parser.__init__` 中新增 `model_config` 参数, 并根据配置设置 `tool_call_id_type`。
3. 清理 `serves` 层 `vllm/entrypoints/openai/chat_completion/serving.py`: 移除原来在 `chat_completion_stream_generator` 和 `chat_completion_full_generator` 中手动计算 `history_tool_call_cnt` 和调用 `make_tool_call_id` 的逻辑, 改为将 `model_config` 传入 `parser`, 由 `parser` 内部自动管理。
4. 同步清理 `vllm/entrypoints/openai/responses/` 相关文件: 在 `context.py`、`utils.py`、`serving.py` 中删除已不再需要的 `get_tool_call_id_type` 引用和手动设置逻辑。
5. 补充测试覆盖: 在 `tests/parser/test_parse.py` 中新增 5 个测试用例, 覆盖非流式解析时 Kimi K2 ID 格式、历史 tool call 偏移、随机 ID deferred 等场景; 在 `tests/parser/test_streaming.py` 中新增 2 个流式测试用例; 删除不再合适的 e2e 测试 (

test_completion_with_function_calling.py 中的 test_tool_id_kimi_k2) , 因为逻辑已下沉到 parser。

关键文件:

- vllm/parser/utils.py (模块 解析器; 类别 source; 类型 core-logic; 符号 count_tool_calls, count_chat_history_tool_calls, count_response_history_tool_calls, count_history_tool_calls) : 新增文件, 提供历史 tool call 计数的核心工具函数, 是整个 ID 生成逻辑的基础。
- vllm/parser/abstract_parser.py (模块 解析器; 类别 source; 类型 core-logic; 符号 _initialize_history_tool_call_cnt, _make_tool_call_id) : 核心逻辑变更: 新增 _initialize_history_tool_call_cnt 和 _make_tool_call_id 方法, 在 parser 内部管理 ID 生成。
- tests/parser/test_parse.py (模块 测试; 类别 test; 类型 test-coverage; 符号 make_parser, test_parse_required_tool_choice_kimi_k2_ids, test_parse_required_tool_choice_kimi_k2_ids_after_history, test_count_history_tool_calls_responses_request) : 新增 5 个针对性测试用例, 覆盖 Kimi K2 ID 格式、历史偏移、随机 ID deferred 等场景, 确保逻辑正确性。
- vllm/entrypoints/openai/chat_completion/serving.py (模块 服务层; 类别 source; 类型 core-logic) : 移除 serving 层原来手动计算历史计数和 ID 生成的逻辑 (+17/-68) , 改为将 model_config 传入 parser 由内部处理, 是架构简化的关键。
- tests/parser/test_streaming.py (模块 测试; 类别 test; 类型 test-coverage; 符号 make_parser, test_parse_delta_required_tool_choice_kimi_k2_ids, test_parse_delta_required_tool_choice_kimi_k2_ids_after_history) : 新增 2 个流式测试用例, 验证流式场景下 ID 生成正确性。
- tests/entrypoints/openai/chat_completion/test_completion_with_function_calling.py (模块 测试; 类别 test; 类型 test-coverage; 符号 k2_server, k2_client, test_tool_id_kimi_k2) : 删除不再合适的 e2e 测试 test_tool_id_kimi_k2, 因为 ID 生成逻辑已下沉到 parser, 应由 parser 级别测试覆盖。
- vllm/parser/engine/parser_engine.py (模块 解析器; 类别 source; 类型 dependency-wiring) : 需要传递 model_config 参数给 Parser 构造函数, 配合架构修改。
- vllm/entrypoints/openai/responses/context.py (模块 服务层; 类别 source; 类型 core-logic) : 移除 get_tool_call_id_type 引用, 清理不再使用的导入。
- vllm/entrypoints/openai/responses/utils.py (模块 服务层; 类别 source; 类型 core-logic) : 移除 get_tool_call_id_type 引用, 与 Responses API 侧清理同步。
- vllm/entrypoints/openai/responses/serving.py (模块 服务层; 类别 source; 类型 core-logic) : 移除 get_tool_call_id_type 引用, 与 Responses API 侧清理同步。
- tests/entrypoints/openai/responses/test_parsable_context_unit.py (模块 测试; 类别 test; 类型 test-coverage) : 配合 Responses 侧清理, 调整测试以移除不再存在的导入依赖。

关键符号: count_tool_calls, count_chat_history_tool_calls, count_response_history_tool_calls, count_history_tool_calls,

```
_initialize_history_tool_call_cnt, _make_tool_call_id, test_parse_required_tool_choice_kimi_k2_ids, test_parse_required_tool_choice_kimi_k2_ids_after_history, test_parse_delta_required_tool_choice_kimi_k2_ids, test_parse_delta_required_tool_choice_kimi_k2_ids_after_history
```

关键源码片段

vllm/parser/utils.py

新增文件，提供历史 tool call 计数的核心工具函数，是整个 ID 生成逻辑的基础。

```
# vllm/parser/utils.py
# SPDX-License-Identifier: Apache-2.0
# SPDX-FileCopyrightText: Copyright contributors to the vLLM project

from collections.abc import Iterable, Sequence
from openai.types.responses import ResponseFunctionToolCall
from vllm.entrypoints.chat_utils import ChatCompletionMessageParam
from vllm.entrypoints.openai.chat_completion.protocol import ChatCompletionRequest
from vllm.entrypoints.openai.responses.protocol import (
    ResponseInputOutputItem,
    ResponsesRequest,
)

def count_tool_calls(tool_calls: object) -> int:
    # 处理 None、字符串、字节、字典或可迭代对象，统一返回工具调用数量
    if tool_calls is None:
        return 0
    if isinstance(tool_calls, (str, bytes, dict)):
        return 1
    if isinstance(tool_calls, Iterable):
        return sum(1 for _ in tool_calls)
    return 1

def count_chat_history_tool_calls(
    messages: Sequence[ChatCompletionMessageParam],
) -> int:
    # 遍历 ChatCompletion 消息列表，统计所有 assistant 角色携带的 tool_calls 总数
    return sum(
        count_tool_calls(msg.get("tool_calls"))
        for msg in messages
        if isinstance(msg, dict) and msg.get("role") == "assistant"
    )

def count_response_history_tool_calls(
    response_items: Sequence[ResponseInputOutputItem],
```

```

) -> int:
    # 遍历 Responses API 的 input items, 支持 ResponseFunctionToolCall 对象和字典格式
    count = 0
    for item in response_items:
        if isinstance(item, ResponseFunctionToolCall):
            count += 1
            continue
        if isinstance(item, dict):
            item_type = item.get("type")
            if item_type == "function_call":
                count += 1
            elif item.get("role") == "assistant":
                count += count_tool_calls(item.get("tool_calls"))
    return count

```

```

def count_history_tool_calls(
    request: ChatCompletionRequest | ResponsesRequest,
) -> int:
    # 根据请求类型分发到不同的统计函数, 统一对外接口
    if isinstance(request, ChatCompletionRequest):
        return count_chat_history_tool_calls(request.messages)
    request_input = request.input
    if isinstance(request_input, str):
        return 0
    return count_response_history_tool_calls(request_input)

```

vllm/parser/abstract_parser.py

核心逻辑变更: 新增 `_initialize_history_tool_call_cnt` 和 `_make_tool_call_id` 方法, 在 parser 内部管理 ID 生成。

```

# vllm/parser/abstract_parser.py 中新增的关键方法
from vllm.entrypoints.chat_utils import (
    get_tool_call_id_type,
    make_tool_call_id,
)
from vllm.parser.utils import count_history_tool_calls

# StreamState 数据类新增字段
@dataclass
class StreamState:
    # ... 原有字段 ...
    history_tool_call_cnt: int = 0
    history_tool_call_cnt_initialized: bool = False # 新增: 防止重复初始化
    tool_call_id_type: str = "random"
    # ...

# Parser.__init__ 中新增 model_config 参数
class Parser:

```

```

def __init__(
    self,
    tokenizer: TokenizerLike,
    tools: list[Tool] | None = None,
    *args,
    model_config=None, # 新增：模型配置，用于决定 ID 类型
    **kwargs,
):
    # ...
    self._stream_state = StreamState(
        tool_call_id_type=(
            get_tool_call_id_type(model_config) # 根据模型配置设置 ID 类型
            if model_config is not None
            else "random"
        ),
        engine_based=self._engine_based,
    )

```

```

def _initialize_history_tool_call_cnt(
    self,
    request: ChatCompletionRequest | ResponsesRequest,
) -> None:
    """
    根据请求历史初始化 tool call 计数器。
    仅当 ID 类型为 kimi_k2 时执行实际统计，
    避免不必要的性能开销。
    """

    state = self._stream_state
    if state.history_tool_call_cnt_initialized:
        return
    if state.tool_call_id_type != "kimi_k2":
        state.history_tool_call_cnt_initialized = True
        return
    state.history_tool_call_cnt = count_history_tool_calls(request)
    state.history_tool_call_cnt_initialized = True

```

```

def _make_tool_call_id(self, function_name: str) -> str | None:
    """
    生成 tool call ID。
    对于 kimi_k2 类型，返回格式为 functions.<function_name>:<idx> 的 ID，
    并自增计数器；其他类型返回 None 表示由上游逻辑处理。
    """

    state = self._stream_state
    if state.tool_call_id_type != "kimi_k2":
        return None
    tool_call_id = make_tool_call_id(
        id_type=state.tool_call_id_type,
        func_name=function_name,
        idx=state.history_tool_call_cnt,
    )

```

```
)  
state.history_tool_call_cnt += 1  
return tool_call_id
```

评论区精华

仅有一条评论来自作者 [chaunceyjiang](#)，解释删除 e2e 测试的原因："Since `call_id` has been moved from serving into the parser, this e2e test is no longer appropriate." 表明团队认同将 ID 生成逻辑下沉后，原有的端到端测试已过时，应该由 parser 级别的单元测试替代。

- 删除 e2e 测试的原因 (testing): 同意删除该 e2e 测试，由 parser 级别的单元测试覆盖。

风险与影响

- 风险：回归风险：serving 层删除了大量手动 ID 生成逻辑，如果 parser 中的新逻辑未正确触发或历史计数初始化失败，可能导致无 ID 或 ID 重复。不过新增的单元测试和流式测试覆盖了主要路径（包括无历史、有历史、流式三种场景），降低了回归概率。兼容性风险：Kimi K2 的 ID 格式从可能为空的 ID 变为固定格式 `functions.<name>:<idx>`，依赖原始空 ID 的下游代码可能受影响。但该格式与 PR 期望一致，属于 bug fix 而非 breaking change。Responses API 影响：`count_response_history_tool_calls` 新增函数涉及 Responses API 的历史统计，若 Responses 请求结构发生变化（如新增类型）可能未同步更新，但当前测试覆盖了基本路径。
- 影响：用户视角：使用 Kimi K2 模型并设置 `tool_choice="required"` 的用户会获得正确的 tool call ID，且流式场景下 ID 也会正确传递。其他模型的工具调用行为不变。系统视角：ID 生成逻辑从 serving 层统一集中到 parser 中，降低了后续新增模型或 ID 格式时的维护成本。新增的 `vllm/parser/utils.py` 成为可复用的工具模块。团队协作：架构更清晰——parser 负责解析和 ID 生成，serving 仅负责编排。
- 风险标记：回归风险：serving 层大量删除可能遗漏，测试覆盖充分但 e2e 测试删除，仅 Kimi K2 模型行为改变

关联脉络

- PR #46314 [Frontend] Port seed_oss to the streaming parser engine as a Qwen3 subclass: 同为 parser 层重构，将特定模型逻辑下沉到 parser 中，与本 PR 的 "逻辑内聚" 方向一致。
- PR #46583 [Rust Frontend] Introduce unified parser interface & combined parser: 统一解析器接口，与本 PR 共同推动 parser 层职责清晰化。