

PR #42748 完整报告

vllm-project/vllm

[Bugfix] Expose usage field in GenerateResponse for disaggregated serving

合并时间: 2026-07-01 18:00

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/42748>

执行摘要

- 一句话: 暴露 disagg 响应中 usage 等字段
- 推荐动作: 值得精读的样板级小 bugfix: 展示了如何通过 Pydantic 模型定义缺少字段导致数据丢失, 以及 AI 辅助审查能快速发现连带字段缺失和参数名不匹配等隐蔽问题。对维护 disagg 相关代码的开发者尤其有参考价值。

功能与动机

disaggregated serving 场景下, 客户端需要获取 token 用量信息 (尤其是 `prompt_tokens_details.cached_tokens`) 以监控缓存命中率并优化前缀复用。但 `GenerateResponse` 缺少 `usage` 字段定义, 导致服务端已计算的值被 Pydantic 丢弃, `model` 和 `created` 也面临同样问题。PR body 明确写道: 'The `usage` field was missing from the `GenerateResponse` model, so pydantic silently dropped it.' Issue 评论中用户 @codeByJiaDong 提供了详细复现证据, 确认该 bug 影响 vLLM 0.20.0。

实现拆解

1. 在 `protocol.py` 中补充字段定义: 为 `GenerateResponse` 类添加 `model: str | None = None`、`created: int | None = None` 和 `usage: UsageInfo | None = Field(default=None)` 三个可选字段。其中 `usage` 使用 `Field(default=None)` 保持与 `GenerateStreamResponse` 一致。
2. 修复构造函数调用的参数名: 在 `serving.py` 的 `serve_tokens_full_generator` 方法中, 将构建 `GenerateResponse` 时的 `id=request_id` 改为 `request_id=request_id`, 确保实际 request ID 被正确传入, 而非使用模型默认的随机 UUID。
3. 对齐流式与非流式响应: 流式响应 `GenerateStreamResponse` 早已包含 `usage` 字段, 本次改动使非流式 `GenerateResponse` 与之对等, `model` 和 `created` 也是 `GenerateStreamResponse` 已有的字段。

关键文件:

- `vllm/entrypoints/scale_out/token_in_token_out/protocol.py` (模块入口层; 类别 source; 类型 core-logic; 符号 `GenerateResponse`): 核心修复所在: 为 `GenerateResponse` 类新增 `usage`、`model`、`created` 三个可选字段, 使其与 `GenerateStreamResponse` 对等, 并防止 Pydantic 静默丢弃数据。

- vllm/entrypoints/scale_out/token_in_token_out/serving.py (模块入口层; 类别 source; 类型 core-logic; 符号 serve_tokens_full_generator) : 修复构造函数调用: 将 id=request_id 改为 request_id=request_id, 确保正确传递请求 ID 而非生成随机值。

关键符号: serve_tokens_full_generator

关键源码片段

vllm/entrypoints/scale_out/token_in_token_out/protocol.py

核心修复所在: 为 GenerateResponse 类新增 usage、model、created 三个可选字段, 使其与 GenerateStreamResponse 对等, 并防止 Pydantic 静默丢弃数据。

```
# vllm/entrypoints/scale_out/token_in_token_out/protocol.py

class GenerateResponse(BaseModel):
    request_id: str = Field(
        default_factory=lambda: f"{random_uuid()}",
        description=(
            "The request_id related to this request. If the caller does "
            "not set it, a random_uuid will be generated. This id is used "
            "through out the inference process and return in response."
        ),
    )
    # 新增: model 与 created 之前被构造时传入但被 Pydantic 静默丢弃
    model: str | None = None
    created: int | None = None
    choices: list[GenerateResponseChoice]
    # 新增: 与 GenerateStreamResponse 保持一致的 usage 字段
    # 确保 prompt_tokens_details.cached_tokens 等信息能被客户端接收
    usage: UsageInfo | None = Field(default=None)
    prompt_logprobs: list[dict[int, Logprob] | None] | None = None
    kv_transfer_params: dict[str, Any] | None = Field(
        default=None,
        description="KVTransfer parameters used for disaggregated serving.",
    )
```

评论区精华

Gemini Code Assist 自动审查指出两个问题: 一是 `GenerateResponse` 除了 `usage` 外还缺少 `model` 和 `created` 字段; 二是 `serving.py` 中传参 `id=request_id` 与模型字段名 `request_id` 不匹配, 导致实际 request ID 被随机 UUID 替代。这两点均在后续提交中修复。此外无人工争议, reviewer @NickLucche 最终批准合并。

- `GenerateResponse` 缺少 `model` 和 `created` 字段 (correctness): 作者随后补充了 `model` 和 `created` 字段到 `GenerateResponse` 中。
- `id=request_id` 参数名不匹配 (correctness): 作者将 `serving.py` 中的 `id=request_id` 修正为 `request_id=request_id`。

风险与影响

- 风险：风险极低：新增字段均为 Optional 并带 None 默认值，不会破坏现有调用方；仅修正构造函数传参名，不改变任何计算逻辑；变更仅涉及 2 个文件共 4 行新增、2 行删除。唯一的边缘情况是，如果外部代码显式检查 GenerateResponse 的字段集合（如 `__fields__`），新增字段可能引发意外行为，但此类使用极少且属于良性变更。
- 影响：直接影响：disaggregated serving 场景下，使用 GenerateResponse 的客户端现在能正确获得 usage（含 `prompt_tokens_details.cached_tokens`）、model 和 created 字段，同时 request_id 不再被错误替换为随机值。影响范围限于 `vllm/entrypoints/scale_out/token_in_token_out/` 模块下的非流式生成响应，不涉及其他 API 端点或核心推理逻辑。
- 风险标记：暂无

关联脉络

- 暂无明显关联 PR