

PR #50816 完整报告

vllm-project/vllm

[Frontend] Require cache_salt to be non-empty via schema

合并时间: 2026-08-03 14:51

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/50816>

执行摘要

- 一句话: cache_salt 非空约束改为 schema 内 min_length, 并移除重复验证器
- 推荐动作: 建议结合 #50764 一起阅读, 理解从 "运行时验证" 到 "schema 验证" 的迁移动机。重点检查 pooling 协议是否应该补上 min_length, 并确认 OpenAI 端点测试是否仍覆盖空字符串拒绝场景。值得学习的小型重构范式, 但需要在合并前补齐遗漏。

功能与动机

PR body 明确指出: "Instead of using model validators, set the min_length directly in the Pydantic field so that the requirement is specified in the generated OpenAPI schema." 这是对 #50764 的 follow-up, 解决其 review 讨论 (#discussion_r3701158217) 中提出的问题——原验证器逻辑只在运行时生效, 客户端无法从 OpenAPI 文档预知 cache_salt 必须非空, 导致间歇性 500 错误难以排查。

实现拆解

1. 在 4 个 OpenAI 协议文件 (chat_completion、completion、responses、scale_out/token_in_token_out) 的 cache_salt Field 中添加 min_length=1, 让 Pydantic 原生执行非空校验并输出到 JSON Schema。
2. 删除 chat_completion、completion、responses 三个文件中重复的 check_cache_salt_support 验证器 (@model_validator mode="before"), 因为其功能已被 min_length 覆盖。
3. 删除 pooling/base/protocol.py 中的 check_cache_salt_support 验证器, 但该文件未同步添加 min_length, 造成约束缺口。
4. 未新增或修改测试文件, 现有测试依赖运行时行为, 是否覆盖 schema 级校验存疑。

关键文件:

- vllm/entrypoints/openai/chat_completion/protocol.py (模块 协议层; 类别 source; 类型 core-logic; 符号 check_cache_salt_support): 主要 OpenAI 聊天端点, cache_salt 字段新增 min_length=1 并删除 check_cache_salt_support 验证器, 代表核心变更模式。
- vllm/entrypoints/pooling/base/protocol.py (模块 协议层; 类别 source; 类型 core-logic; 符号 check_cache_salt_support): 只删除了验证器但未添加 min_length, 导致 Pooling 端点 cache_salt 非空约束丢失, 是本 PR 中最值得关注的风险点。

- `vllm/entrypoints/openai/completion/protocol.py` (模块 协议层; 类别 source; 类型 core-logic; 符号 `check_cache_salt_support`) : 补全接口的 `cache_salt` 字段同样添加 `min_length=1` 并删除验证器, 保持行为一致。
- `vllm/entrypoints/openai/responses/protocol.py` (模块 协议层; 类别 source; 类型 core-logic; 符号 `check_cache_salt_support`) : Responses API 的 `cache_salt` 字段添加 `min_length=1`, 删除验证器, 是 OpenAI 系列协议的一部分。
- `vllm/entrypoints/scale_out/token_in_token_out/protocol.py` (模块 协议层; 类别 source ; 类型 configuration) : 该协议原先没有非空验证, 现在新增 `min_length=1`, 属于行为收紧, 需注意兼容性。

关键符号: `check_cache_salt_support`

关键源码片段

`vllm/entrypoints/openai/chat_completion/protocol.py`

主要 OpenAI 聊天端点, `cache_salt` 字段新增 `min_length=1` 并删除 `check_cache_salt_support` 验证器, 代表核心变更模式。

```
class ChatCompletionRequest(OpenAIBaseModel):
    # ... 其他字段 ...

    cache_salt: str | None = Field(
        default=None,
        # 关键变更: min_length=1 让非空约束直接写进 OpenAPI schema,
        # 客户端可在调用前发现该要求; 原来仅存在于 model_validator 中。
        min_length=1,
        description=(
            "If specified, the prefix cache will be salted with the provided "
            "string to prevent an attacker to guess prompts in multi-user "
            "environments. The salt should be random, protected from "
            "access by 3rd parties, and long enough to be "
            "unpredictable (e.g., 43 characters base64-encoded, corresponding "
            "to 256 bit).",
        ),
    )

    # 注意: 此前定义的 check_cache_salt_support (@model_validator mode="before")
    # 已在本 PR 中删除, 因为 min_length 已覆盖其“非空字符串”校验逻辑;
    # 这样避免了自定义错误路径与 schema 声明不一致的问题。
```

`vllm/entrypoints/pooling/base/protocol.py`

只删除了验证器但未添加 `min_length`, 导致 Pooling 端点 `cache_salt` 非空约束丢失, 是本 PR 中最值得关注的风险点。

```
class PoolingBasicRequestMixin(OpenAIBaseModel):
    # ... 其他字段 ...

    cache_salt: str | None = Field(
```

```
default=None,
description=(
    "If specified, the prefix cache will be salted with the provided "
    "string to prevent an attacker to guess prompts in multi-user "
    "environments. The salt should be random, protected from "
    "access by 3rd parties, and long enough to be "
    "unpredictable (e.g., 43 characters base64-encoded, corresponding "
    "to 256 bit).",
),
)
# 注意：此处没有添加 min_length=1，却删除了原有的
# check_cache_salt_support 验证器 —— 相比其他协议，这里丢掉了
# cache_salt 非空约束，可能导致空字符串被静默接受，疑似本 PR 遗漏。
```

评论区精华

PR 无实质 review 评论，AndreasKaratzas 直接批准。claude[bot] 因 fork 来源自动禁用 review。因此没有围绕设计权衡的公开讨论。

- 暂无高价值评论线程

风险与影响

- 风险：
 1. pooling/base/protocol.py 删除了验证器但未补 min_length，导致 Pooling 端点 cache_salt 空字符串会静默通过，与原行为不一致，属于回归风险。
 2. chat_completion 等文件错误消息从自定义 VLLMValidationError（含 parameter 字段）变为 Pydantic 默认错误（String should have at least 1 character），依赖旧错误格式的客户端或测试可能受影响。
 3. scale_out/token_in_token_out 协议原先没有非空校验，现在新增 min_length 是行为收紧，可能影响之前传入空字符串的调用。
 4. 无新增测试，schema 层面的校验行为未被显式验证。 - 影响：影响所有使用 cache_salt 的 OpenAI 兼容端点（chat、completion、responses、scale_out）以及 pooling 端点。正面影响是 OpenAPI 文档更准确，客户端可在调用前发现约束；负面影响是 pooling 约束丢失和可能的错误消息变化。对内部实现减少了约 50 行重复验证代码，维护性提升。整体影响范围中等，但回归点明确。 - 风险标记：Pooling 端非空约束丢失，缺少测试覆盖，错误消息格式变化，scale_out 行为收紧

关联脉络

- PR #50764 [Bugfix][Frontend] Constrain Anthropic cache_salt to non-empty: 前一 PR 通过 model_validator 实现非空约束，本 PR 是其 follow-up，将约束迁移到 schema 层以解决 OpenAPI 文档缺口。
- PR #49498 [Frontend] Add cache_salt support to Anthropic Messages API: 引入 cache_salt 字段，本 PR 完善了该字段的 schema 级校验约束。