

PR #51260 完整报告

vllm-project/vllm

[Bugfix] Skip fetching revision for model when model and weights_model are different

合并时间: 2026-08-07 20:55

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/51260>

执行摘要

- 一句话: 修复权重与模型仓库不同时 revision 误解析
- 推荐动作: 值得快速阅读, 修复方案简洁且根因分析清晰。可以关注 `ModelConfig.__post_init__` 中多仓库 sourcing 的边界条件处理方式, 以及如何用 mock 测试避免真实网络调用。该 PR 展示了配置文件中对多来源 (model/config/weights/tokenizer) 进行差异化处理的典型模式。

功能与动机

PR #49990 之后, nightly CI 上两个 GGUF 测试 (`plugins_tests/gguf/test_gguf_plugin_generate.py::test_models` 的两个用例) 出现 `huggingface_hub RevisionNotFoundError`。根因是 `ModelConfig` 可以分别从不同 repo 加载 config 和 weights, 而 revision 是共享字段, `resolve_revision` 被无条件作用于 `self.model`, 导致 weights repo (如 `unsloth/Qwen3-0.6B-GGUF`) 收到 base repo (`Qwen/Qwen3-0.6B`) 的 commit hash, `snapshot_download` 校验时 404。此前 revision 为 `None` 时 `huggingface_hub` 会自行解析 main 分支, 反而正确。

实现拆解

1. 修改 revision 解析前置条件: 在 `vllm/config/model.py` 的 `__post_init__` 中, 将原先仅判断 `hf_config_path` 是否与 model 同源的 `can_resolve_model_revision`, 扩展为同时要求 `model_weights` 为空或与 model 相等 (`weights_from_model`), 只有 config 与 weights 都来自 model 时才对 `self.model` 调用 `resolve_revision`。
2. 保留 tokenizer 分支逻辑: `tokenizer_revision` 的解析不受影响, 仍按原逻辑独立处理, 避免引入额外行为变化。
3. 新增回归测试: 在 `tests/test_config.py` 中新增 `test_revision_not_resolved_when_weights_differ_from_model` 和 `test_revision_resolved_when_weights_match_model`, 分别验证 weights 与 model 不同源时 revision 保持 `None`、同源时正确解析为 `ResolvedRevision`, 并通过 mock `resolve_revision` 断言调用行为。
4. 配套调整: 测试文件新增 `from huggingface_hub import ResolvedRevision` 导入, 并顺带优化了 `vllm.config.vllm` 的导入写法。

关键文件:

- `vllm/config/model.py` (模块 模型配置; 类别 source; 类型 data-contract) : 核心修改文件, 调整 `ModelConfig` 初始化时 `revision` 解析的条件, 避免 `weights repo` 与 `model repo` 不同时错误解析 `shared revision`。
- `tests/test_config.py` (模块 模型配置; 类别 test; 类型 test-coverage; 符号 `test_revision_not_resolved_when_weights_differ_from_model`, `test_revision_resolved_when_weights_match_model`) : 新增两个回归测试, 覆盖 `weights` 与 `model` 同源 / 异源两种情形, 防止该 bug 复发; 测试通过 `mock resolve_revision` 避免了真实网络请求。

关键符号: `ModelConfig.post_init`, `resolve_revision`

关键源码片段

`vllm/config/model.py`

核心修改文件, 调整 `ModelConfig` 初始化时 `revision` 解析的条件, 避免 `weights repo` 与 `model repo` 不同时错误解析 `shared revision`。

`vllm/config/model.py` `ModelConfig.__post_init__` 中 `revision` 解析逻辑 (head 版本)

```
self.model = maybe_model_redirect(self.model)
# 默认 tokenizer 与 model 一致
if self.tokenizer is None:
    self.tokenizer = self.model
if self.tokenizer_revision is None:
    self.tokenizer_revision = self.revision
self.tokenizer = maybe_model_redirect(self.tokenizer)

if isinstance(self.hf_config_path, str):
    self.hf_config_path = maybe_model_redirect(self.hf_config_path)

# ... hf_overrides 处理 ...

self.maybe_pull_model_tokenizer_for_runai(self.model, self.tokenizer)

# 关键修复: 权重来自不同 repo 时, 不能基于 self.model 预先解析 revision
# 否则会把 base repo 的 commit hash 传给 weights repo, 导致 404
weights_from_model = not self.model_weights or self.model_weights == self.model
# 配置来自不同 repo 同理, 也无法用 self.model 解析
config_from_model = not self.hf_config_path or self.hf_config_path == self.model
can_resolve_model_revision = config_from_model and weights_from_model

if can_resolve_model_revision:
    # 仅在 config 与 weights 都源自 model 时, 提前解析并缓存 revision,
    # 避免下游重复解析; 否则保持 None, 由 huggingface_hub 按各 repo 自行解析
    self.revision = resolve_revision(
        self.model,
        self.revision,
        self.hf_token,
```

```

)

# tokenizer_revision 走独立逻辑，不受上述条件限制
if (
    can_resolve_model_revision
    and self.tokenizer == self.model
    and self.tokenizer_revision == requested_revision
):
    self.tokenizer_revision = self.revision
else:
    self.tokenizer_revision = resolve_revision(
        self.tokenizer,
        self.tokenizer_revision,
        self.hf_token,
    )

```

tests/test_config.py

新增两个回归测试，覆盖 weights 与 model 同源 / 异源两种情形，防止该 bug 复发；测试通过 mock resolve_revision 避免了真实网络请求。

```

# tests/test_config.py 新增回归测试 (head 版本)

from huggingface_hub import ResolvedRevision

# 使用真实存在的 Qwen3-0.6B revision 作为测试值
REVISION = "c1899de289a04d12100db370d81485cdf75e47ca"

@patch("vllm.config.model.resolve_revision", return_value=ResolvedRevision(REVISION))
def test_revision_not_resolved_when_weights_differ_from_model(mock_resolve):
    # model 与 weights 来自不同 repo (GGUF 场景)，revision 不应被提前解析
    model_weights = "unsloth/Qwen3-0.6B-GGUF:Q8_0"
    config = ModelConfig("Qwen/Qwen3-0.6B", model_weights=model_weights)
    assert config.revision is None

@patch("vllm.config.model.resolve_revision", return_value=ResolvedRevision(REVISION))
def test_revision_resolved_when_weights_match_model(mock_resolve):
    # 默认情况下 weights 来自 model 自身，revision 应被正常解析
    model = "Qwen/Qwen3-0.6B"
    config = ModelConfig(model)
    assert isinstance(config.revision, ResolvedRevision)
    assert config.revision.resolved == REVISION
    mock_resolve.assert_any_call(model, None, config.hf_token)

```

评论区精华

该 PR review 评论极少，无实质技术争论。hmellor（合入者）直接 APPROVED，并帮助提交了 clarify、merge main 与 fix test 等后续提交；claude[bot] 因 PR 来自 fork 而自动跳过审

查。身体中作者详细给出了根因分析和测试计划，结论清晰。

- 暂无高价值评论线程

风险与影响

- 风险：

1. 回归风险：can_resolve_model_revision 条件收紧后，当 model_weights 与 model 不同源时，revision 不再被预先解析。若后续代码（如 snapshot_download 之外）依赖 config.revision 已解析，可能拿到 None；但 PR 中 tokenizer_revision 分支仍独立解析，且 weights 下载路径原本就应使用原始 revision，因此风险较低。
2. 兼容性：model_weights 支持带子路径或格式后缀（如 repo:Q8_0），目前判断仅做字符串全等比较，若用户以不同写法（如末尾斜杠、大小写）传同一 repo，可能误判为不同源而跳过预解析，但实际下载仍会由 huggingface_hub 正确处理，影响有限。
3. 测试覆盖：新增测试仅覆盖 GGUF 场景的两种分支，未覆盖 hf_config_path 与 model_weights 同时存在且与 model 不同的组合，但现有逻辑已覆盖。- 影响：影响范围：修复 GGUF 模型（或其他 model 与 weights 分仓库的场景）加载时的 revision 解析错误，消除 nightly CI 中的软失败。对普通用户而言，之前无法加载的跨仓库组合现在可正常工作；对维护者，CI 稳定性提升。影响程度为小范围 bugfix，核心路径（配置初始化）有改动但逻辑简单。- 风险标记：核心配置路径变更，需关注 revision 为 None 的兼容性，修复缺少更多组合场景测试

关联脉络

- PR #49990 [Bugfix] (IMPORTANT) Resolve HF revision once to avoid repeated network calls: PR #49990 引入了对 model revision 的提前解析，导致 model 与 weights 异源时 weights repo 收到错误的 commit hash，本 PR 是对其的回归修复。