

PR #48873 完整报告

vllm-project/vllm

[CI] Extend max-model-len for `test\_parsable\_context` to allow reasoning to finish

合并时间: 2026-07-17 11:36

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/48873>

执行摘要

- 一句话: 延长测试模型长度上限, 修复推理不完整
- 推荐动作: 可作为临时修复合入, 但建议尽快跟进分析推理 token 用量波动的原因, 并考虑是否完全移除 max_model_len 限制以彻底解决 flaky 问题。

功能与动机

CI 中 test_parsable_context.py::test_basic[Qwen/Qwen3-8B] 频繁报错: response.status 为 'incomplete', 原因是 max_model_len 设为 5000 时不足以容纳推理所需的全部 token, 导致响应被截断。PR body 附有 CUDA 和 ROCm 上的 flaky 失败链接。

实现拆解

修改测试 fixture 中 server 启动参数, 将 max_model_len 从 5000 调整为 6000。

- 文件: tests/entrypoints/openai/responses/test_parsable_context.py 第 41 行
- 仅变更一个数字, +1/-1
- 该值仍为任意选择的宽松值, 作者提到可考虑完全移除该参数, 由 vLLM 自动选择默认值

关键文件:

- tests/entrypoints/openai/responses/test_parsable_context.py (模块测试; 类别 test; 类型 test-coverage): 唯一变更文件, 将 max_model_len 从 5000 改为 6000 以稳定 flaky 测试

关键符号: 未识别

关键源码片段

tests/entrypoints/openai/responses/test_parsable_context.py

唯一变更文件, 将 max_model_len 从 5000 改为 6000 以稳定 flaky 测试

```
# tests/entrypoints/openai/responses/test_parsable_context.py (line 37-49)
@pytest.fixture(scope="module")
def server():
    args = [
        "--reasoning-parser",
        "qwen3",
```

```
    "--max_model_len",
    "6000", # 从 5000 增加到 6000，确保推理能完整输出（有时需约 5423 tokens）
    "--structured-outputs-config.backend",
    "xgrammar",
    "--enable-auto-tool-choice",
    "--tool-call-parser",
    "hermes",
    "--tool-server",
    "demo",
]
```

评论区精华

审核人 AndreasKaratzas 同意作为快速稳定方案，但指出需在后续 PR 中调查 token 用量的非确定性波动。

- 临时修复与后续调查 (other): 同意当前 PR 作为临时修复，后续需调查波动根源。

风险与影响

- 风险：风险极低：仅修改测试配置参数，不影响任何生产代码或测试逻辑。增大 `max_model_len` 可能增加单次测试内存占用，但 6000 对 Qwen3-8B 影响可忽略。
- 影响：仅影响 CI 中 `test_parsable_context.py::test_basic` 测试的稳定性，使其在 ROCm 和 CUDA 上更可靠地通过。无其他影响。
- 风险标记：临时修复

关联脉络

- 暂无明显关联 PR