

PR #52282 完整报告

vllm-project/vllm

[CI] Harden RemoteVLLMServer GPU cleanup checks

合并时间: 2026-08-20 03:17

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/52282>

执行摘要

- 一句话: RemoteVLLMServer 关闭超时改用引擎清理宽限, 稳定 CI
- 推荐动作: 建议 CI/ 测试基础设施维护者精读 tests/utils.py 中 `_get_process_termination_timeout` 的设计与单测断言 ($75 = 60 + 15$), 理解测试等待如何与引擎清理宽限对齐; 一般读者快速浏览即可。值得关注的设计决策: 把引擎超时语义以钩子形式暴露给测试辅助类, 使 vLLM 运行时关闭策略的变更能自动传导到测试基建。

功能与动机

PR body 明确列出三项要求: 在 RemoteVLLMServer 中保留解析出的请求关闭超时; 对 RemoteOpenAIServer 将 15 秒外层宽限叠加到 `get_engine_process_shutdown_timeout` 返回的 EngineCore 进程超时之上; 并用派生超时同时驱动 `subprocess.Popen.wait` 与自定义 multiprocessing 服务器的 `Process.join`。结合测试侧改动可见, 旧测试 `test_abort_timeout_exits_quickly` 硬编码 4 秒等待, 而引擎关闭仍可能占用平台相关的资源清理宽限 (如 #52281 中引入的 15 秒), 导致 GPU/ROCm CI 偶发超时失败, 因此需要让测试等待与引擎清理语义对齐。

实现拆解

1. 基类超时钩子 (tests/utils.py): 新增 `_get_process_termination_timeout()` 返回 15.0; `__init__` 中保存 `self._request_shutdown_timeout = float(args.shutdown_timeout)`; `_terminate_process_tree` 的 `proc.wait(timeout=15)` 改为调用该钩子, 使通用远程服务器保留原有 15 秒终止截止时间。
2. 子类叠加引擎宽限 (tests/utils.py): RemoteOpenAIServer 重写钩子, 调用 `get_engine_process_shutdown_timeout(self._request_shutdown_timeout, self._request_shutdown_timeout)`, 断言非 None 后加基类 15 秒; RemoteOpenAIServerCustom 的 `proc.join(15)` 同样改为调用钩子, 保证自定义多进程服务器也遵循派生超时。
3. 单元测试 (tests/entrypoints/unit_tests/test_remote_vllm_server.py 新增): 通过 monkeypatch 验证 `_terminate_process_tree` 将请求超时 (0.0, 0.0) 原样传给引擎超时计算, 且最终 `wait(timeout=75.0)` (60 秒引擎超时 + 15 秒外层宽限)。
4. 集成测试调整 (tests/entrypoints/launchers/test_shutdown.py): `test_abort_timeout_exits_quickly` 重命名为 `test_abort_timeout_exits_within_cleanup_grace`, 移除 4 秒硬编码等待与耗时断言, 改用 `remote_server._get_process_termination_t`

imeout() 作为超时，失败信息也更明确地指向资源清理宽限。

5. 格式收尾 (vllm/platforms/cpu.py) : 仅拆分超过行长限制的注释与 logger 调用，无行为变化，属于 pre-commit 格式修复。

关键文件：

- tests/utils.py (模块 测试工具；类别 test；类型 core-logic；符号 `_get_process_termination_timeout`, `_terminate_process_tree`, `RemoteOpenAIServer`, `RemoteOpenAIServerCustom`) : 承载核心逻辑：新增 `_get_process_termination_timeout`、保存请求关闭超时并统一 `_terminate_process_tree` 的等待时间，是测试基建的关键改动。
- tests/entrypoints/unit_tests/test_remote_vllm_server.py (模块 服务器测试；类别 test；类型 test-coverage；符号 `test_openai_server_shutdown_wait_covers_engine_cleanup`, `get_engine_timeout`) : 新增单测验证 `_terminate_process_tree` 将请求超时传给引擎超时计算并最终使用 75 秒等待，是本 PR 行为的关键回归护栏。
- tests/entrypoints/launchers/test_shutdown.py (模块 关闭测试；类别 test；类型 test-coverage；符号 `test_abort_timeout_exits_within_cleanup_grace`) : 将快速退出断言改为宽限内退出断言，消除硬编码 4 秒超时导致的 CI 偶发失败。
- vllm/platforms/cpu.py (模块 平台层；类别 source；类型 cleanup) : 仅 pre-commit 格式修复（注释与 logger 调用换行），无逻辑变化，体现 PR 收尾的格式清理。

关键符号： `_get_process_termination_timeout`, `_terminate_process_tree`, `test_openai_server_shutdown_wait_covers_engine_cleanup`, `test_abort_timeout_exits_within_cleanup_grace`, `get_engine_timeout`

关键源码片段

tests/utils.py

承载核心逻辑：新增 `_get_process_termination_timeout`、保存请求关闭超时并统一 `_terminate_process_tree` 的等待时间，是测试基建的关键改动。

tests/utils.py (整理后片段，非原始 diff 摘录)

```
from vllm.v1.engine.utils import get_engine_process_shutdown_timeout
```

```
class RemoteVLLMServer:
```

```
    # 基类保留 15 秒通用终止宽限，与新引入的引擎清理宽限保持一致
```

```
    def _get_process_termination_timeout(self) -> float:
```

```
        return 15.0
```

```
    def __init__(self, ...):
```

```
        ...
```

```
        # 保存解析出的请求关闭超时，供派生超时计算使用
```

```
        self._request_shutdown_timeout = float(args.shutdown_timeout)
```

```
        ...
```

```
    def _terminate_process_tree(self) -> None:
```

```
        ...
```

```

print(f"[RemoteOpenAIServer] Sent SIGTERM to process {pid}")
try:
    # 等待时间 = 引擎清理宽限 + 15 秒外层宽限, 不再硬编码 15
    self.proc.wait(timeout=self._get_process_termination_timeout())
except subprocess.TimeoutExpired:
    # Phase 2: SIGKILL 整个进程组
    ...

class RemoteOpenAIServer(RemoteVLLMServer):
    # OpenAI 服务器叠加引擎关闭超时: 请求超时同时作为 wait 与 engine 参数,
    # 返回 None 时显式断言失败, 避免静默回退到短超时
    def _get_process_termination_timeout(self) -> float:
        engine_timeout = get_engine_process_shutdown_timeout(
            self._request_shutdown_timeout,
            self._request_shutdown_timeout,
        )
        assert engine_timeout is not None
        return engine_timeout + super()._get_process_termination_timeout()

    def _terminate_process_tree(self) -> None:
        ...
        self.proc.terminate()
        # 自定义多进程服务器同样使用派生超时
        self.proc.join(self._get_process_termination_timeout())

```

tests/entrypoints/unit_tests/test_remote_vllm_server.py

新增单测验证 `_terminate_process_tree` 将请求超时传给引擎超时计算并最终使用 75 秒等待, 是本 PR 行为的关键回归护栏。

```

# tests/entrypoints/unit_tests/test_remote_vllm_server.py
from unittest.mock import Mock

import pytest

import tests.utils as test_utils
from tests.utils import RemoteOpenAIServer

def test_openai_server_shutdown_wait_covers_engine_cleanup(
    monkeypatch: pytest.MonkeyPatch,
):
    # 直接构造对象, 避免真正拉起子进程
    server = object.__new__(RemoteOpenAIServer)
    server._request_shutdown_timeout = 0.0
    server.proc = Mock(pid=1234)
    engine_timeout_args = []

    def get_engine_timeout(request_timeout, process_timeout):
        engine_timeout_args.append((request_timeout, process_timeout))

```

```

return 60.0

monkeypatch.setattr(
    test_utils, "get_engine_process_shutdown_timeout", get_engine_timeout
)
monkeypatch.setattr(test_utils.os, "getpgid", lambda pid: pid)
monkeypatch.setattr(server, "_kill_process_group_survivors", Mock())

server._terminate_process_tree()

# 请求超时 0 应原样传给引擎超时计算
assert engine_timeout_args == [(0.0, 0.0)]
# 60 秒引擎超时 + 15 秒外层宽限 = 75 秒
server.proc.wait.assert_called_once_with(timeout=75.0)

```

tests/entrypoints/launchers/test_shutdown.py

将快速退出断言改为宽限内退出断言，消除硬编码 4 秒超时导致的 CI 偶发失败。

```

# tests/entrypoints/launchers/test_shutdown.py
@pytest.mark.asyncio
@pytest.mark.parametrize("wait_for_engine_idle", [0.0, 2.0])
async def test_abort_timeout_exits_within_cleanup_grace(
    wait_for_engine_idle: float,
):
    server_args = [
        "--dtype", "bfloat16",
        "--max-model-len", "256",
        "--enforce-eager",
        "--gpu-memory-utilization", "0.05",
        "--max-num-seqs", "4",
        "--shutdown-timeout", "0",
    ]

    with RemoteOpenAIServer(MODEL_NAME, server_args) as remote_server:
        proc = remote_server.proc
        child_pids = _get_child_pids(proc.pid)

        if wait_for_engine_idle > 0:
            ...
        proc.send_signal(signal.SIGTERM)

        # 请求超时为 0 会立即中止请求，但进程回收仍可能占用平台
        # 特定的资源清理宽限，因此等待时间改用派生超时
        termination_timeout = remote_server._get_process_termination_timeout()
        try:
            proc.wait(timeout=termination_timeout)
        except subprocess.TimeoutExpired:
            proc.kill()
            proc.wait(timeout=5)

```

```

    pytest.fail(
        "Process did not exit within the resource cleanup grace "
        f"({termination_timeout}s)"
    )

    assert proc.returncode in (0, -15, None), f"Unexpected: {proc.returncode}"
    await _assert_children_cleaned_up(child_pids)

```

评论区精华

PR 没有实质 review 评论。claude[bot] 仅自动提示该仓库配置了手动 review；

DarkLight1337 给出 APPROVED（空评论）。提交历史可见 deadline 语义的多轮演进：先是引入超时钩子与请求超时保存，随后依次提交 'Fix RemoteVLLMServer pre-commit errors'、'Simplify RemoteVLLMServer shutdown grace'、'Fix abort shutdown cleanup deadline'、'Fix CPU pre-commit formatting'，说明关闭截止时间计算与格式经过了多轮迭代才稳定。

- 无实质 review 评论，DarkLight1337 直接批准 (other)：在 5 个提交与多轮 CI（含 AMD nightly）后合入。

风险与影响

- 风险：
 1. 测试等待时长可能显著增加：RemoteOpenAIServer 关闭等待从固定 15 秒变为引擎超时（默认 60 秒级别）加 15 秒（单测中为 75 秒），失败路径下测试总时长变长。
 2. 语义耦合：tests/utils.py 的关闭逻辑依赖 vllm.v1.engine.utils.get_engine_process_shutdown_timeout 的返回值语义，若上游策略调整（如返回 None 或更短超时），会直接影响所有使用 RemoteVLLMServer 的测试；当前以 assert 强制非 None，属于显式失败而非优雅降级。
 3. 回归风险低：生产代码仅 cpu.py 格式调整，但 tests/utils.py 是共享测试基建，改动会影响大量基于 RemoteOpenAIServer 的测试的关闭路径，需要关注 CI 整体耗时与偶发稳定性。
 - 影响：对生产用户无影响（纯测试 /CI 变更）。对团队而言：消除 abort timeout (--shutdown-timeout 0) 场景下 GPU/ROCm CI 的偶发关闭超时失败，减少重跑成本；统一了测试辅助类的关闭等待语义，并为平台特定超时提供了可被子类覆盖的 _get_process_termination_timeout 钩子。测试基建的语义与 v1 EngineCore 关闭策略保持一致，未来调整引擎宽限时测试可自动跟随。
 - 风险标记：测试等待时长可能延长，依赖引擎关闭超时语义，assert 硬失败于 None 超时

关联脉络

- PR #52281 [ROCm] Give EngineCore cleanup grace after request abort: 本 PR 的 15 秒清理宽限与 get_engine_process_shutdown_timeout 均源自该 PR，是同一关闭语义的测试侧配套。
- PR #52981 [CI/Build] Fix CPU platform pre-commit formatting: 与本 PR 最后一个提交的 CPU pre-commit 格式修复同类，都是 CI 格式收尾工作。