

PR #46069 完整报告

vllm-project/vllm

[CPU][Bugfix][Speculative Decoding] Accept USE_FP64_GUMBEL in CPU recovered-tokens sampler

合并时间: 2026-06-23 19:54

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/46069>

执行摘要

- 一句话: 修复 CPU 推测解码因 USE_FP64_GUMBEL 崩溃
- 推荐动作: 此 PR 是一次必要且安全的回归修复, 改动极小 (+3/-1 行), 建议合入。值得关注的是: 当为 GPU 内核添加新参数时, 需同步更新所有 CPU / 回退封装函数以避免类似 API 不一致错误。

功能与动机

PR #43150 为 GPU Triton 内核 `sample_recovered_tokens_kernel` 及其调用者添加了 `USE_FP64_GUMBEL` 参数, 但未更新 `vllm/utils/cpu_triton_utils.py` 中的 CPU 回退封装函数 `_sample_recovered_tokens_kernel_impl`。导致 CPU 推测解码 (草稿模型 / EAGLE3) 初始化时抛出 `TypeError: _sample_recovered_tokens_kernel_impl() got an unexpected keyword argument 'USE_FP64_GUMBEL'`, EngineCore 启动失败。从 v0.23.0 开始存在回归, v0.22.1 不受影响。

实现拆解

1. 添加参数声明: 在 `vllm/utils/cpu_triton_utils.py` 的 `_sample_recovered_tokens_kernel_impl` 函数签名末尾添加 `USE_FP64_GUMBEL=False`, 与 GPU Triton 内核签名对齐。
2. 转换数据类型: 在底层 C++ 内核调用前, 将 `inv_q` 通过 `.to(torch.float32)` 显式转换为 `float32`, 因为 CPU C++ 内核通过 `data_ptr<float>()` 读取 `inv_q`, 始终使用 `fp32` 精度。
3. 代码注释: 在类型转换处添加注释 `# C++ kernel reads inv_q as float32`. 以明确设计意图。
4. 影响范围: 仅影响 CPU 构建 (`VLLM_TARGET_DEVICE=cpu`) 下的推测解码回退路径, GPU Triton 内核路径保持不变。

关键文件:

- `vllm/utils/cpu_triton_utils.py` (模块 CPU 工具; 类别 `source`; 类型 `core-logic`; 符号 `_sample_recovered_tokens_kernel_impl`): 这是唯一修改的文件, 添加了 `USE_FP64_GUMBEL` 参数并修正 `inv_q` 数据类型以匹配底层 C++ 内核。

关键符号: `_sample_recovered_tokens_kernel_impl`

关键源码片段

vllm/utils/cpu_triton_utils.py

这是唯一修改的文件，添加了 `USE_FP64_GUMBEL` 参数并修正 `inv_q` 数据类型以匹配底层 C++ 内核。

vllm/utils/cpu_triton_utils.py 中修改后的函数

```
def _sample_recovered_tokens_kernel_impl(
    output_token_ids,
    cu_num_draft_tokens,
    draft_token_ids,
    draft_probs,
    target_probs,
    inv_q,
    vocab_size,
    BLOCK_SIZE=None,
    NO_DRAFT_PROBS=False,
    USE_FP64_GUMBEL=False, # 新增参数，与 GPU Triton 内核 API 对齐
):
    # C++ reads integer tensors as int64_t*; ensure correct dtype.
    orig_dtype = output_token_ids.dtype
    output_i64 = _ensure_int64(output_token_ids)
    torch.ops._C.sample_recovered_tokens_kernel_impl(
        output_i64,
        _ensure_int64(cu_num_draft_tokens),
        _ensure_int64(draft_token_ids),
        draft_probs,
        target_probs,
        # C++ kernel reads inv_q as float32.
        inv_q.to(torch.float32), # 显式转换，因为底层始终使用 fp32 采样
        vocab_size,
        NO_DRAFT_PROBS,
    )
    if orig_dtype != torch.int64:
        output_token_ids.copy_(output_i64.to(orig_dtype))
```

评论区精华

PR 无 review 评论，仅有一条自动机器人欢迎消息和作者提交的 CI 运行请求 (`/gcrun`)。作者在评论中说明了 CI 失败 (ROCm 和 Docker Hub 问题) 与本次 CPU 专用改动无关。

- 暂无高价值评论线程

风险与影响

- 风险：风险极低。变更仅限 CPU 回退封装函数的签名和内部类型转换，GPU 路径不受影响。新增参数默认为 `False`，对常见路径无行为变化。类型转换 `inv_q.to(torch.float32)` 是安全的，因为底层 C++ 内核始终以 `float*` 读取。未提供自动化测试覆盖 CPU 推测解码路径，但 PR 描述中作者的手动测试验证了功能正常。

- 影响：影响范围：仅修复 CPU 推测解码（draft model / EAGLE3）的初始化崩溃回归。影响程度：高，因为该崩溃完全阻塞了 CPU 上推测解码的使用。非推测解码的普通 CPU 服务不受影响。
- 风险标记：缺少直接测试覆盖，回归修复（v0.23.0 引入）

关联脉络

- PR #43150 Fix FP64 Gumbel precision coverage: 本 PR 修复了 #43150 引入的回归，后者为 GPU 内核添加了 USE_FP64_GUMBEL 但未更新 CPU 回退函数。
- PR #37798 FP64 Gumbel for speculative decoding: #43150 的原始基础，首次引入 fp64 Gumbel 精度改进。