

PR #47383 完整报告

vllm-project/vllm

[Bugfix][Model Runner V2][Spec Decode] Fix int32 offset overflow in block verification kernels

合并时间: 2026-07-02 22:32

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/47383>

执行摘要

- 一句话: 修复 V2 投机解码 block 验证 kernel int32 索引溢出
- 推荐动作: 值得精读, 展示 Triton 内核中因混合 int32/int64 导致的索引溢出典型修复模式, 测试设计 (参数化溢出场景) 值得学习。

功能与动机

当词汇表规模约为 155k 时, `logit_idx * vocab` 等乘法在 int32 范围内溢出, 导致 block 验证内核读取越界 logits, 进而破坏接受决策或引起非法内存访问。此问题在审计 V2 spec-decode 采样路径时发现 (关联 #47239), 虽非该 issue 直接原因, 但属于同一类 int32 索引溢出 bug。

实现拆解

1. 在 `_compute_cumulative_log_p_kernel` 中, 将 `req_state_idx` 和 `start_idx` 的加载结果通过 `.to(tl.int64)` 提升为 int64。
2. 在 `_compute_local_residual_mass_kernel` 中, 将 `logit_idx` (来自 `tl.program_id(0)`) 和 `req_state_idx` 也通过 `.to(tl.int64)` 提升。
3. 新增测试文件 `tests/v1/worker/test_gpu_rejection_sampler_i64.py`, 构造两种溢出场景 (目标端索引 15000、草稿端索引 8000), 验证修复后的输出与无溢出参考一致。
4. 每个测试用例分配约 5 GiB GPU 内存, 并在 CI 中仅对 CUDA 运行。

关键文件:

- `vllm/v1/worker/gpu/spec_decode/rejection_sampler_utils.py` (模块 投机解码; 类别 source; 类型 core-logic; 符号 `_compute_cumulative_log_p_kernel`, `_compute_local_residual_mass_kernel`): 核心修复文件, 修改两个 Triton kernel 中的 int32 索引提升
- `tests/v1/worker/test_gpu_rejection_sampler_i64.py` (模块 测试; 类别 test; 类型 test-coverage; 符号 `_run`, `test_block_verification_i64_indexing`): 新增回归测试, 覆盖目标端和草稿端索引溢出场景

关键符号: `_compute_cumulative_log_p_kernel`, `_compute_local_residual_mass_kernel`

评论区精华

PR 无实质性 review 讨论，仅由 claude[bot] 触发审查提示和 TheEpicDolphin 批准。

- 无实质性 Review 讨论 (other): 直接合并

风险与影响

- 风险：风险较低，修复直接且被测试覆盖。但需注意其他未检查的 Triton 内核是否也存在类似问题。性能影响可忽略，int64 加载在 GPU 上通常无开销。
- 影响：影响使用 rejection_sample_method='block' 的投机解码用户，尤其大词汇表模型（如 GLM）。修复后输出正确，避免崩溃或静默错误。对其他路径无影响。
- 风险标记：int32 溢出，非默认路径，GPU 内核修复

关联脉络

- PR #46560 [Bugfix][Sampler] Fix int32 overflow in sampler kernels: 与本 PR 修复同一类 int32 索引溢出 bug，在采样器内核中已修复
- PR #47239 [Audit] Spec-decode sampling path audit: 审计过程中发现此类溢出问题，但此 PR 并非直接原因