

PR #46560 完整报告

vllm-project/vllm

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

合并时间: 2026-06-25 02:00

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/46560>

执行摘要

- 一句话: 修复 MRV2 采样器 int32 偏移溢出
- 推荐动作: 此 PR 为高价值修复, 值得特别关注。展示了在 Triton kernel 中防范 int32 溢出的模式, 建议团队学习并纳入编码规范。

功能与动机

DFlash (以及任何推测多个 token 的方法) 在 V1 引擎启动时崩溃, 报错 'RuntimeError: Triton Error [CUDA]: an illegal memory access was encountered'。根因是 MRV2 采样器和拒绝采样 Triton kernel 使用 int32 计算每行偏移, 当 $\text{batch_size} * \text{vocab_size}$ 超过 2^{31} 时发生溢出 (warmup 使用 963 个请求, K=15 时 $96316=15408$ 行, $15408151936 \approx 2.34e9 > 2.147e9$)。

实现拆解

1. 识别所有受影响 Triton kernel: 涉及 topk_topp_triton、logit_bias、penalties、min_p、gumbel、bad_words、logprob 及 rejection_sampler_utils 中的共 13 个 kernel。
2. 在每个 kernel 入口, 将 `tl.program_id(0)` 和从指针加载的索引 (如 `expanded_idx_mapping`) 显式转换为 `tl.int64`, 确保后续偏移计算 (如 `token_idx * logits_stride`) 不会在 int32 下溢出。
3. 添加两个回归测试: `test_large_batch_int64_row_offset` 构造 $\text{batch_size} * \text{vocab_size} > 2^{31}$ 的输入, 验证最高偏移行的 top-k 结果与首行一致; `test_token_logprobs_large_batch_int64_row_offset` 验证 log-softmax 结果正确。
4. 验证 DFlash 贪婪输出与无推测解码的贪婪输出 bit 一致, 且所有现有单元测试通过, 微基准性能无变化。

关键文件:

- `vllm/v1/worker/gpu/spec_decode/rejection_sampler_utils.py` (模块 推测解码; 类别 source; 类型 core-logic; 符号 `_compute_block_stats_kernel`, `_rejection_kernel`, `_resample_kernel`): 包含三个核心 rejection kernel 的修复, 是溢出的直接触发点。
- `tests/v1/sample/test_topk_topp_sampler.py` (模块 采样器测试; 类别 test; 类型 test-coverage; 符号 `test_large_batch_int64_row_offset`): 新增回归测试, 直接验证 int64 修复的正确性。

- vllm/v1/worker/gpu/sample/gumbel.py (模块 Gumbel 采样; 类别 source; 类型 core-logic; 符号 `_temperature_kernel`, `gumbel_block_argmax`, `_gumbel_sample_kernel`) : Gumbel 采样 kernel 也受溢出影响, 修改了三个函数。
- tests/v1/sample/test_logprobs.py (模块 logprobs 测试; 类别 test; 类型 test-coverage ; 符号 `test_token_logprobs_large_batch_int64_row_offset`) : 新增 logprobs 路径的回归测试。
- vllm/v1/worker/gpu/sample/logprob.py (模块 logprobs 计算; 类别 source; 类型 core-logic; 符号 `_topk_log_softmax_kernel`, `_ranks_kernel`) : logprob 计算 kernel 也受溢出影响。
- vllm/v1/worker/gpu/sample/bad_words.py (模块 禁用词过滤; 类别 source; 类型 core-logic; 符号 `_bad_words_kernel`) : bad_words kernel 也受溢出影响。
- vllm/v1/worker/gpu/sample/logit_bias.py (模块 logit 偏置; 类别 source; 类型 core-logic; 符号 `_bias_kernel`) : logit_bias kernel 也受溢出影响。
- vllm/v1/worker/gpu/sample/min_p.py (模块 MinP 采样; 类别 source; 类型 core-logic; 符号 `_min_p_kernel`) : min_p kernel 也受溢出影响。
- vllm/v1/worker/gpu/sample/penalties.py (模块 惩罚项; 类别 source; 类型 core-logic; 符号 `_penalties_kernel`) : penalties kernel 也受溢出影响。
- vllm/v1/sample/ops/topk_topp_triton.py (模块 TopKTopP 采样; 类别 infra; 类型 infrastructure) : TopKTopP Triton kernel 是主修复目标之一。

关键符号: `_compute_block_stats_kernel`, `_rejection_kernel`, `_resample_kernel`, `_temperature_kernel`, `gumbel_block_argmax`, `_gumbel_sample_kernel`, `_topk_log_softmax_kernel`, `_ranks_kernel`, `_bad_words_kernel`, `_bias_kernel`, `_min_p_kernel`, `_penalties_kernel`, `_topk_topp_kernel`

关键源码片段

vllm/v1/worker/gpu/spec_decode/rejection_sampler_utils.py

包含三个核心 rejection kernel 的修复, 是溢出的直接触发点。

```
# 以 _compute_block_stats_kernel 为例, 展示入口处类型提升
@triton.jit
def _compute_block_stats_kernel(
    target_local_argmax_ptr, target_local_argmax_stride,
    target_local_max_ptr, target_local_max_stride,
    target_local_sumexp_ptr, target_local_sumexp_stride,
    draft_local_max_ptr, draft_local_max_stride,
    draft_local_sumexp_ptr, draft_local_sumexp_stride,
    target_logits_ptr, target_logits_stride,
    draft_logits_ptr, draft_logits_stride_0, draft_logits_stride_1,
    expanded_idx_mapping_ptr, expanded_local_pos_ptr,
    temp_ptr, vocab_size, num_speculative_steps,
    BLOCK_SIZE: tl.constexpr, HAS_DRAFT_LOGITS: tl.constexpr,
):
    logit_idx = tl.program_id(0).to(tl.int64) # 转换为 int64, 避免后续偏移计算 (如 logit_idx * logits_
```

```

stride) 溢出
draft_step_idx = tl.load(expanded_local_pos_ptr + logit_idx)

if draft_step_idx >= num_speculative_steps:
    return

req_state_idx = tl.load(expanded_idx_mapping_ptr + logit_idx).to(tl.int64) # 同样转换为 int64

```

tests/v1/sample/test_topk_topp_sampler.py

新增回归测试，直接验证 int64 修复的正确性。

```

@large_gpu_mark(min_gb=24)
def test_large_batch_int64_row_offset(self):
    """Regression: per-row offset (row * vocab_size) must not overflow int32.

    Speculative decoding expands the logits batch (e.g. DFlash drafts K
    tokens per request), so batch_size * vocab_size can exceed 2**31. With
    int32 offset arithmetic the per-row pointer wraps to a negative address
    and the kernel hits a CUDA illegal memory access. Use a batch where
    batch_size * vocab_size > 2**31 and give the highest-offset row the same
    logits as row 0: an overflow there would read a different row and change
    the kept set.
    """
    from vllm.v1.sample.ops.topk_topp_triton import apply_top_k_top_p_triton

    if not current_platform.is_cuda():
        pytest.skip("int32 row-offset overflow is a CUDA kernel issue")
    vocab_size = 131072
    batch_size = 2**31 // vocab_size + 64 # batch_size * vocab_size > 2**31
    required_bytes = batch_size * vocab_size * 4 + (1 << 30)
    if torch.cuda.mem_get_info()[0] < required_bytes:
        pytest.skip(f"needs ~{required_bytes / 1e9:.0f} GB of free GPU memory")

    logits = torch.randn(batch_size, vocab_size, generator=self.generator, dtype=torch.float32)
    logits[batch_size - 1] = logits[0] # 最高偏移行与第 0 行相同
    k = torch.full((batch_size,), 5, dtype=torch.int32)
    result = apply_top_k_top_p_triton(logits, k, None)
    torch.accelerator.synchronize() # 让异步非法内存访问立即浮现
    kept_first = (result[0] > float("-inf")).nonzero(as_tuple=True)[0]
    kept_last = (result[batch_size - 1] > float("-inf")).nonzero(as_tuple=True)[0]
    assert kept_first.numel() == 5, f"row 0 kept {kept_first.numel()}, expected 5"
    assert torch.equal(kept_first, kept_last), (
        "highest-offset row produced a different top-k mask than the "
        "identical row 0 (int32 row-offset overflow)"
    )

```

评论区精华

审核人 WoosukKwon 认为“虽然部分更改可能并非严格必要，但从防御性编程角度看是合理的”（感谢修复）。无其他讨论线程。

- 防御性编程与必要性 (design): 接受更改，认为从防御性编程角度合理。

风险与影响

- 风险：变更为纯类型提升，GPU 全局地址本就是 64 位，性能无影响。风险主要在于：若后续在其他平台（如 ROCm）上 Triton 行为有差异，可能需要验证。但修改仅在 CUDA 路径，且测试已覆盖。
- 影响：修复 Model Runner V2 路径下推测解码在任意 batch 时的启动崩溃和运行时非法内存访问。对非推测解码用户无影响，因 batch 远小于溢出阈值。
- 风险标记：涉及 10 个文件，CUDA kernel 变更，影响推测解码路径

关联脉络

- 暂无明显关联 PR