

PR #34462 完整报告

sgl-project/sglang

[Triton] Bound the sliding-window extend-attention KV loop: -86.6% on SWA layers, -9.4% prefill GPU, bit-identical

合并时间: 2026-08-24 15:58

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/34462>

执行摘要

- 一句话: 限制滑动窗口 KV 循环下界, SWA prefill 提速 86.6%
- 推荐动作: 值得精读。该 PR 是 Triton 内核循环界优化的教科书式案例: 用 mask 条件反推 tile 下界、利用 constexpr 折叠保证非目标分支零指令差异, 并以位级一致性、正反 A/B 顺序、编译产物对比和跨厂商复现构建了完整验证链。对 SGLang 内核维护者, 建议把这种验证方法论沉淀为 SWA 内核修改的常规流程; 后续若有人改动 SKIP_TILE 或循环边界, 应先跑 PR 中记录的 91 组覆盖作为回归基线。

功能与动机

extend 阶段 KV 循环的 causal 上界按 m-block 计算, 但下界被硬编码为 0, 导致滑动窗口层会遍历 chunk 里每一个更早的 64 宽 tile, 并依赖 `SKIP_TILE = tl.max(tl.max(final_mask, 1), 0) == 0` 来跳过 MMA, 每次浪费都会触发跨 wave 的 `tl.max` 归约。作者在 PR body 中给出的量化证据是: 滑动窗口层耗时 507 μ s, 仅比全注意力层的 518 μ s 便宜 2.1%, 尽管它理论上只需要 192 个 tile 而不是 4160 个——结构性优势几乎全部消耗在遍历与掩码上。

实现拆解

1. 变更入口: 唯一修改文件是 `python/sglang/kernels/ops/attention/extend_attention.py`, 改动集中在 `_fwd_kernel` 的 extend 阶段 KV 循环 (+10/-1)。
2. 核心逻辑: 原代码是 `for start_n in range(0, extend_end, BLOCK_N)`, 新代码先计算 `extend_start`。窗口条件 `q <= kv + SLIDING_WINDOW_SIZE` 决定了 `m0` 之前不可能存在未掩码元素, 因此可贡献的最小 `BLOCK_N` 对齐 tile 是 `(maximum(m0 - W, 0) // BLOCK_N) * BLOCK_N`, 其中 `m0 = cur_block_m * BLOCK_M`。
3. 分支与折叠设计: `SLIDING_WINDOW_SIZE` 是 `tl.constexpr`, 当 `W <= 0` 时新的 if 分支在 trace 期被折叠, 循环退回原始行为, 全注意力层的 AMDGCN 产物仅 DWARF 元数据不同, 1305 条指令零差异。SKIP_EXTEND 为真时 `extend_end = 0`, range 为空, 行为不变。
4. 配套与测试: 本 PR 未新增或修改单测, 正确性以 PR body 中的数值实验代替: 91 组配置 (窗口 1-4096、`W > seq_len`、ragged batch、前后缀、fp16/bf16、多种 head_dim 与 tiling) `max_abs` 恰为 0.0, 并用独立 fp32 torch 参考确认两边都在算真实的 windowed attention; 另有 H200 独立复现。

关键文件：

- python/sglang/kernels/ops/attention/extend_attention.py（模块 注意力内核；类别 source；类型 core-logic；符号 _fwd_kernel）：唯一变更文件，修改 _fwd_kernel 的 extend 阶段 KV 循环下界，是全部性能收益与位级一致性声明的来源。

关键符号：_fwd_kernel

关键源码片段

python/sglang/kernels/ops/attention/extend_attention.py

唯一变更文件，修改 _fwd_kernel 的 extend 阶段 KV 循环下界，是全部性能收益与位级一致性声明的来源。

```
# _fwd_kernel 内，extend 阶段的 KV 循环（片段，已重整）
else:
    cur_seq_len_extend = tl.minimum(
        cur_seq_len_extend, (cur_block_m + 1) * BLOCK_M
    )
extend_end = 0 if SKIP_EXTEND else cur_block_m_end

# 滑动窗口下界推导：mask 条件为 q <= kv + SLIDING_WINDOW_SIZE,
# 因此 m0 = cur_block_m * BLOCK_M 之前不可能有未掩码元素；
# extend_start 取可贡献的最小 BLOCK_N 对齐 tile。
# SLIDING_WINDOW_SIZE 是 tl.constexpr，W <= 0 时该分支在 trace 期折叠，
# 全注意力层的循环边界与改动前逐字一致，编译产物零指令差异。
extend_start = 0
if SLIDING_WINDOW_SIZE > 0:
    extend_start = (
        tl.maximum(cur_block_m * BLOCK_M - SLIDING_WINDOW_SIZE, 0) // BLOCK_N
    ) * BLOCK_N

for start_n in range(extend_start, extend_end, BLOCK_N):
    start_n = tl.multiple_of(start_n, BLOCK_N)
    mask_n = (start_n + offs_n) < cur_block_m_end
    # 原 SKIP_TILE 判定仍然保留：对于窗口紧贴 m0 的对角 tile，
    # 若整块无有效元素仍会被跳过，但不再需要为纯浪费的远处 tile 做跨 wave 归约
```

评论区精华

该 PR 没有任何 review 评论（review_comments_count = 0）。唯一的外部讨论是 Hert4 在 issue 评论区给出的 NVIDIA H200 独立验证：

Independent corroboration on NVIDIA H200 (SM90). This PR is tagged [ROCm] and all the numbers in the description are gfx950, but the change is architecture independent...

Hert4 用 head_dim = 256、window = 1024 的 Gemma 4 滑窗层直接测 extend_attention_fwd，seq_len 4096 到 65536 得到 1.35x 到 8.22x 的速度提升，控制了无滑动窗口的 head_dim = 512 作为对照，确认改动只影响窗口层。这条评论补足了非 ROCm

硬件的数据点，也侧面验证了作者对 constexpr 折叠与架构无关性的论断。

- H200 独立复现验证架构无关性 (testing): 验证通过：优化与架构无关，收益随 seq_len 增长而放大；该数据与 PR body 中 ROCm 侧结论相互印证。

风险与影响

- 风险:

1. 正确性风险：加速依赖一个推导出的 tile 下界公式，作者声明对每个 residue class 独立验证过且不要求 $BLOCK_M \geq BLOCK_N$ ，但仓库没有配套单测；91 组配置虽覆盖面广，仍未穷尽所有后端 tiling 与未来 Triton 行为。
2. 编译期折叠风险：全注意力层的零指令差异依赖 `SLIDING_WINDOW_SIZE` 为 `tl.constexpr` 且编译器在 trace 期内联展开 if 分支；若未来把该值改为运行时参数或 Triton 版本改变常量传播行为，需要重新评估。
3. 性能回归面：改动只影响 extend 阶段循环下界，decode 路径与全注意力层不受影响；1k/1k 短 chunk 负载下 p50 波动在 $\pm 2.6\%$ 内、无趋势，判定为中性而非回退。
4. 维护风险：公式中 `// BLOCK_N` 后再乘回 `BLOCK_N` 的对齐写法依赖 reader 对整数运算的理解，代码注释已说明原理，但未来改动时容易被误简化。
 - 影响：影响用户面：所有使用 Triton attention 后端、且模型含滑动窗口注意力层（如 gpt-oss-120b、Gemma 4 类）的 prefill 请求都将显著收益；对无 SWA 的模型是无副作用的零变化。服务指标上，gpt-oss-120b TP4 ISL 8192 场景 total prefill GPU 从 551.6 ms 降到 499.5 ms (-9.4%)，`_fwd_kernel` 占比从 24.18% 降到 15.39%；run_sweep 显示 8k/1k 负载随并发提升单调改善（p50 -2.4% 到 -5.5%，p99 最高 -11.9%），TPOT 持平或更好。影响程度属于中高价值、低风险的性能优化，团队侧收益集中在 AMD gfx950 与 NVIDIA 等各架构的 prefill 吞吐与首 token 时延。
 - 风险标记：位级一致性依赖 constexpr 折叠，缺少新增单测覆盖，编译产物差异依赖 Triton 常量传播，公式对齐行为需维护者复核

关联脉络

- PR #34461 [Triton]（本 PR 的 stacked 前置 PR）：PR body 明确声明 Stacked on #34461，该分支的第一个 commit 来自 #34461，作者要求先 review/merge #34461；本 PR 的独立 diff 仅为第二个 commit，即这里分析的 SWA 循环下界优化。