

PR #28392 完整报告

sgl-project/sglang

[AMD] Annotate ATOM source for imported v4 unified attention kernels

合并时间: 2026-06-16 13:12

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/28392>

执行摘要

- 一句话: 为 AMD V4 内核添加 ATOM 来源注释
- 推荐动作: 可作为合规性参考, 无精读必要, 但体现了 AMD 团队对上游代码归属的重视。

功能与动机

PR body 明确指出 "The unified kv attention kernel is porting from ATOM's sparse attention kernels", 添加来源注释以明确 ATOM 的原始出处, 提高代码溯源性和合规性。

实现拆解

在以下 4 个文件开头添加注释行, 标注内核来源为 ATOM 项目:

1. `paged_decode.py`: 添加 `# The following kernel is imported from ATOM.` 和 `# Source: atom/model_ops/v4_kernels/paged_decode.py`。
2. `paged_decode_indices.py`: 添加类似注释, 指向 ATOM 的对应索引内核。
3. `paged_prefill.py`: 添加类似注释, 指向 ATOM 的预填充内核。
4. `deepseek_v4_memory_pool.py`: 在 `DeepSeekV4UnifiedKVPool` 类定义前添加注释 `# The following kv pool follows ATOM's unified_kv kernel layout.`。所有变更均为纯注释添加, 不影响运行时逻辑。

关键文件:

- `python/sglang/srt/layers/attention/dsv4/unified_kv_kernels/paged_decode.py` (模块 解码内核; 类别 `source`; 类型 `core-logic`): 核心 V4 解码注意力内核, 添加 ATOM 来源注释
- `python/sglang/srt/layers/attention/dsv4/unified_kv_kernels/paged_decode_indices.py` (模块 解码内核; 类别 `source`; 类型 `core-logic`): V4 分页解码索引散列内核, 添加 ATOM 来源注释
- `python/sglang/srt/layers/attention/dsv4/unified_kv_kernels/paged_prefill.py` (模块 预填充内核; 类别 `source`; 类型 `core-logic`): V4 稀疏预填充注意力内核, 添加 ATOM 来源注释
- `python/sglang/srt/mem_cache/deepseek_v4_memory_pool.py` (模块 内存池; 类别 `source`; 类型 `core-logic`): V4 统一 KV pool 类, 添加 ATOM 布局注释

关键符号: 未识别

关键源码片段

`python/sglang/srt/layers/attention/dsv4/unified_kv_kernels/paged_decode.py`

核心 V4 解码注意力内核，添加 ATOM 来源注释

```
# SPDX-License-Identifier: MIT
# Copyright (C) 2024-2026, Advanced Micro Devices, Inc. All rights reserved.

# 以下内核从 ATOM 项目导入
# 来源：atom/model_ops/v4_kernels/paged_decode.py
```

"""稀疏解码注意力，基于 unified KV pool 和 per-token 分页索引。

为 V4 decode + CUDAGraph 设计：替代每次前向生成 kv_flat_sa（其形状依赖 n_committed_per_seq，变化会阻止 CUDAGraph 捕获），使用统一 KV pool 通过分页索引访问，类似 `aiter.mla.mla_decode_fwd` 的 API 风格。

调用者约定：

```
unified_kv: [total_pages, D] BF16 (page_size=1)
    概念上合并了 SWA 环形缓冲区和单层 V4 的压缩器分页缓存。
    槽位 [0, swa_pages) 引用 SWA 条目 (state_slot * win + ring);
    槽位 [swa_pages, ...) 引用压缩 K 条目 (block_id * K_PER_BLOCK + slot_in_block)。
kv_indices: [total_indices] int32 — 每个 token 的扁平槽位列表。
    每个 token 的条目位于 kv_indices[kv_indptr[t] : kv_indptr[t+1]]。
    **所有条目必须为 unified_kv.shape[0] 内的有效槽位 ID。**
生产环境中的 decode 索引构建器 (write_v4_paged_decode_indices) 输出 ragged-packed
索引（无哨兵值）；
CG 填充的 token 通过 indptr[t+1] == indptr[t] 获得零长度切片。内核不再进行 per-iter slot >= 0
哨兵检查。
kv_indptr: [N+1] int32 — 真实前缀和（每个 token 长度可变）。
attn_sink: [H] per-head 可学习 softmax 分母偏置 (V4 特有)。
softmax_scale: float。
```

返回：

```
out: [N, H, D], 与 q 相同数据类型。
```

数值：使用 $\log_2$ 域内的 online-softmax ($qk_scale = softmax_scale * \log_2 E$)，注意力吸收作为虚拟 K 合并。在 fp32 累加容差范围内与 PyTorch 参考实现 (`_sparse_attn_ragged_torch`) 位接近。

"""

评论区精华

无 review 争议；一份自动评论达到日配额，与 PR 无关。

- 暂无高价值评论线程

风险与影响

- 风险：无技术风险。仅添加注释，不涉及逻辑变更。
- 影响：影响极小：改善了代码可维护性和溯源透明度，对用户无感知。
- 风险标记：暂无

关联脉络

- 暂无明显关联 PR