

PR #30792 完整报告

sgl-project/sglang

[Kernel] Migrate DSA + DSV4 attention kernels to sglang.kernels (RFC #29630, Phase 2.5, 5/7)

合并时间: 2026-07-15 11:11

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/30792>

执行摘要

- 一句话: 搬迁 DSA/DSV4 注意力内核至 sglang.kernels
- 推荐动作: 建议关注 DeepSeek 相关内核开发的工程师精读此 PR, 了解新内核包的目录结构和导入约定。值得注意的设计决策: 采用字节相同的移动确保功能零回归; 从混合模块中提取纯内核代码到独立文件, 提高了模块化; 多次合并冲突处理展示了大范围重构中的协作模式。

功能与动机

Phase 2.5 的内核统一计划 (RFC #29630) 旨在将分散在 sglang.srt 各处的内核代码集中到 sglang.kernels 下, 降低维护成本, 提高代码复用和可发现性。此 PR 负责处理 DeepSeek 的 DSA 和 DSV4 注意力内核子树。

实现拆解

1. 创建目标目录结构 ops/attention/dsa/ 和 ops/attention/dsv4/, 将现有的 DSA/DSV4 内核文件完整移动到新位置 (字节相同, git 识别为 R100 重命名)。
2. 从混合模块中提取内核代码到独立文件: 例如, 从 dsa/utils.py 将 dsa_cp_round_robin_split_q_seqs_kernel 提取到 dsa/cp_split.py; 从 dsv4/compressor_v2.py 将 HIP 下的 C128 压缩内核提取到 dsv4/compress_c128_hip.py; 从 dsv4/sparse_prefill_utils.py 将 _build_swa_token_ids_kernel 等提取到 dsv4/sparse_prefill_kernels.py。
3. 重写所有受影响的导入语句, 涉及约 60 个源文件、测试文件和基准文件。替换路径如 from sglang.srt.layers.attention.dsv4.dequant_k_cache 改为 from sglang.kernels.ops.attention.dsv4.dequant_k_cache。
4. 在 ops/attention/__init__.py 中注册迁移后的入口点为 KernelSpec 清单。
5. 修复遗留的旧路径引用, 例如 index_buf_accessor 包在 deepseek_v4_memory_pool 中的导入, 以及测试文件 test_dsa_transform_index.py 的导入路径。
6. 多次合并 main 分支, 解决与同系列 PR (3/7、4/7) 共享 __init__.py 文件时的追加冲突。
7. 运行 test_kernels_namespace.py 和 test_fused_op.py 的 33 个测试验证导入完整性。

关键文件:

- `python/sglang/srt/layers/attention/dsv4/compressor_v2.py` (模块 DSV4; 类别 source; 类型 core-logic; 符号 `_c128_compress_decode_kernel`, `_c128_compress_prefill_write_kernel`, `_c128_compress_prefill_compress_kernel`, `_compress_forward_c128_triton`) : DSV4 压缩器核心逻辑文件, 移除了 275 行 HIP 下的 C128 压缩内核, 改为导入新位置的独立包。同时新增了 `_use_online_compress`、`_extract_positions_from_plan` 和 `_compress_forward_c128_fallback` 等函数, 弥补内核外移后的功能空缺。
- `python/sglang/kernels/ops/attention/dsv4/compress_c128_hip.py` (模块 DSV4; 类别 infra; 类型 infrastructure; 符号 `_c128_compress_decode_kernel`, `_c128_compress_prefill_write_kernel`, `_c128_compress_prefill_compress_kernel`, `_compress_forward_c128_triton`) : 新添加的文件, 包含了从 `compressor_v2.py` 中移出的 HIP 下的 C128 压缩 Triton 内核 (`_c128_compress_decode_kernel`, `_c128_compress_prefill_write_kernel` 等), 是内核迁移的目标位置之一。
- `python/sglang/srt/layers/attention/dsv4/sparse_prefill_utils.py` (模块 DSV4; 类别 source; 类型 core-logic; 符号 `_build_swa_token_ids_kernel`, `_combine_topk_swa_indices_kernel`) : DSV4 稀疏预填充工具文件, 移除了两个 Triton 内核 (`_build_swa_token_ids_kernel`、`_combine_topk_swa_indices_kernel`) 到新包, 并修改了 `dequant_k_cache` 的导入路径。
- `python/sglang/srt/layers/attention/dsv4/compress_hip.py` (模块 DSV4; 类别 source; 类型 core-logic; 符号 `_rms_normalize_kernel`, `rms_normalize_triton`) : HIP 压缩器实现, 移除了 `_rms_normalize_kernel` 和 `rms_normalize_triton` 的本地定义, 改为从新包 `sglang.kernels.ops.attention.dsv4.rms_normalize_hip` 导入, 同时调整了 `fused_compress_triton` 的导入路径。
- `python/sglang/srt/layers/attention/dsv4/indexer.py` (模块 DSV4; 类别 source; 类型 core-logic; 符号 `_fused_scale_kernel`, `fused_scale`) : DSV4 Indexer 文件, 移除了 `_fused_scale_kernel` 和 `fused_scale` 函数的本地定义, 改为从新包导入, 同时更新了 `tilelang_kernel` 的导入路径。
- `python/sglang/srt/layers/attention/dsa/utils.py` (模块 DSA; 类别 source; 类型 core-logic; 符号 `dsa_cp_round_robin_split_q_seqs_kernel`) : DSA 工具文件, 移除了 `dsa_cp_round_robin_split_q_seqs_kernel` 的本地定义, 改为从新包 `sglang.kernels.ops.attention.dsa.cp_split` 导入。

关键符号: `_use_online_compress`, `_extract_positions_from_plan`, `_compress_forward_c128_fallback`, `_c128_compress_decode_kernel`, `_c128_compress_prefill_write_kernel`, `_c128_compress_prefill_compress_kernel`, `rms_normalize_triton`, `_rms_normalize_kernel`, `fused_scale`, `_fused_scale_kernel`, `_build_swa_token_ids_kernel`, `_combine_topk_swa_indices_kernel`, `dsa_cp_round_robin_split_q_seqs_kernel`

关键源码片段

`python/sglang/srt/layers/attention/dsv4/compressor_v2.py`

DSV4 压缩器核心逻辑文件，移除了 275 行 HIP 下的 C128 压缩内核，改为导入新位置的独立包。同时新增了 `_use_online_compress`、`_extract_positions_from_plan` 和 `_compress_forward_c128_fallback` 等函数，弥补内核外移后的功能空缺。

```
# python/sglang/srt/layers/attention/dsv4/compressor_v2.py (head)
```

```
# 移除了 HIP 下的 Triton 内核，改为导入新位置的独立包
```

```
from __future__ import annotations
```

```
from typing import TYPE_CHECKING, List, Literal, Optional, TypeAlias, Union, cast
```

```
import torch
```

```
from sglang.jit_kernel.dsv4 import (
```

```
    CompressorDecodePlan,
```

```
    CompressorPrefillPlan,
```

```
    compress_forward,
```

```
    compress_norm_rope_store,
```

```
)
```

```
from sglang.jit_kernel.utils import is_hip_runtime
```

```
from sglang.srt.environ import envs
```

```
_is_hip = is_hip_runtime()
```

```
# 新增：在线压缩开关，仅用于 C128 压缩比
```

```
def _use_online_compress(compress_ratio: int) -> bool:
```

```
    return compress_ratio == 128 and envs.SGLANG_OPT_USE_ONLINE_COMPRESS.get()
```

```
# 新增：从 plan 中提取 RoPE 位置
```

```
def _extract_positions_from_plan(
```

```
    plan: Union[CompressorDecodePlan, CompressorPrefillPlan],
```

```
    compress_ratio: int,
```

```
) -> torch.Tensor:
```

```
    plan_tensor = plan[1]
```

```
    seq_lens = plan_tensor[:, :4].contiguous().view(torch.int32).squeeze(-1)
```

```
    positions = seq_lens.to(torch.int32) - compress_ratio
```

```
    return positions
```

```
# 新增：C128 压缩的 PyTorch fallback
```

```
def _compress_forward_c128_fallback(
```

```
    kv_score_buffer: torch.Tensor,
```

```
    kv_score_input: torch.Tensor,
```

```
    ape: torch.Tensor,
```

```
    plan: Union[CompressorDecodePlan, CompressorPrefillPlan],
```

```
    head_dim: int,
```

```
) -> torch.Tensor:
```

```
    # ... 实现略，封装了在线 softmax pooling 的纯 PyTorch 版本
```

```
    pass
```

`python/sglang/srt/layers/attention/dsv4/compress_hip.py`

HIP 压缩器实现，移除了 `_rms_normalize_kernel` 和 `rms_normalize_triton` 的本地定义，改为从新包 `sglang.kernels.ops.attention.dsv4.rms_normalize_hip` 导入，同时调整了 `fused_compress_triton` 的导入路径。

```
# python/sglang/srt/layers/attention/dsv4/compress_hip.py (head)
# 移除了 Triton 内核定义，改为导入新包

from __future__ import annotations
import os
from functools import cached_property
from typing import TYPE_CHECKING, Any
import torch
import torch.nn as nn

# 导入其他内核位置保持不变
from sglang.kernels.ops.attention.deepseek_v4_rope import (
    apply_rotary_emb_triton,
    fused_norm_rope_inplace_triton,
    fused_softmax_pool_triton,
)
# fused_compress_triton 导入路径从旧 srt 路径改为新 kernels 路径
from sglang.kernels.ops.attention.dsv4.fused_compress_triton import (
    fused_ape_pool_norm_rope,
)
from sglang.srt.env import envs
from sglang.srt.layers.attention.dsa.dsa_indexer import rotate_activation
from sglang.srt.layers.attention.dsv4.compressor import Compressor as _CompressorBase
from sglang.srt.layers.attention.nsa.nsa_indexer import rotate_activation
# ...

# 新增：从新包导入 rms_normalize_triton
from sglang.kernels.ops.attention.dsv4.rms_normalize_hip import rms_normalize_triton

class DeepseekRefRMSNorm(nn.Module):
    def __init__(self, dim: int, eps: float = 1e-6):
        super().__init__()
        self.dim = dim
        self.eps = eps
        self.weight = nn.Parameter(torch.ones(dim, dtype=torch.float32))

    def forward(self, x: torch.Tensor):
        return rms_normalize_triton(x, self.eps, self.weight)

class CompressorHip(_CompressorBase):
    # ... 其余不变
    pass
```

评论区精华

此 PR 的 Review 评论较少，仅有一条自动配额提示和作者 `/rerun-failed-ci` 指令。没有实质性的设计讨论，因为变更为纯机械搬迁，未引入逻辑变化。技术讨论体现在 commit 记录中：多次合并冲突处理记录了与同系列 PR 共享文件（如 `ops/attention/__init__.py`）时的解决策略。

- 暂无高价值评论线程

风险与影响

- 风险：主要风险来自大规模导入路径重写（涉及约 60 个文件）。虽然移动的内核字节完全相同，但任何遗漏的导入更新都会导致运行时 `ImportError`。此 PR 通过 CI 的 33 个导入测试验证了完整性，但 DSV4 端到端测试仅在 H200/B200 等特定硬件上运行，若 CI 未覆盖则存在回归风险。另外，`compress_c128_hip.py` 包含 HIP 专用 Triton 内核，AMD 运行时未测试可能遗漏问题。
- 影响：对用户：无影响，所有公共 API 保持不变。对开发者：未来导入 DSA/DSV4 内核需使用新路径 `sglang.kernels.ops.attention.dsa / dsv4`。对系统：无功能变化，内核逻辑未改。对团队：统一的内核组织方式降低了维护成本，有利于后续优化和 Bug 修复。
- 风险标记：大面积导入重写，缺少 GPU 覆盖测试，核心路径变更

关联脉络

- PR #30784 [Kernel] Migrate scattered quantization kernels to `sglang.kernels` (RFC #29630, Phase 2.5, 1/7): Phase 2.5 系列的第 1 步，与本 PR（第 5 步）共享迁移计划和目标，共同完成内核统一。
- PR #30786 [Kernel] Migrate scattered MoE kernels to `sglang.kernels` (RFC #29630, Phase 2.5, 2/7): Phase 2.5 系列的第 2 步，共享迁移框架和集成测试。
- PR #30787 [Kernel] Migrate top-level srt/layers stray kernels to `sglang.kernels` (RFC #29630, Phase 2.5, 3/7): Phase 2.5 系列的第 3 步，与本 PR 共享 `ops/attention/init.py` 的追加更改，存在合并冲突。
- PR #30789 [Kernel] Migrate generic attention kernels to `sglang.kernels` (RFC #29630, Phase 2.5, 4/7): Phase 2.5 系列的第 4 步，同样共享 `ops/attention/init.py` 的追加更改，存在合并冲突。
- PR #30793 [Kernel] Migrate linear-attention, MiniMax-sparse and diffusion kernels to `sglang.kernels` (RFC #29630, Phase 2.5, 6/7): Phase 2.5 系列的第 6 步，与本 PR 同属一系列，整体推进内核统一目标。