

PR #26788 完整报告

sgl-project/sglang

[JIT Kernel] DeepSeek-V4 DSA indexer: faster top-k + page-table transform (runtime k <= 2048)

合并时间: 2026-07-06 11:23

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/26788>

执行摘要

- 一句话: DeepSeek-V4 稀疏 attention JIT top-k kernel 运行时优化, 最高 1.5 倍加速
- 推荐动作: 该 PR 是 DeepSeek-V4 推理性能的关键优化, 代码质量高 (测试覆盖 172 个 case)。值得精读:
 - `topk_impl.cuh` 中运行时 top-k 的直方图 + 阈值二阶段算法。
 - `plan_topk_v2` 的 plan-routed 动态分派逻辑。
 - MaxSmem 元编程技巧, 减少静态共享内存选择。

对于使用 DeepSeek-V4 的团队, 建议验证目标 GPU 对 PDL 的支持, 并关注未解决的 PDL 触发问题。

功能与动机

DeepSeek-V4 的 DSA indexer 在选择每个查询的 top-k KV 位置并映射到 page table 时, 操作处于解码关键路径, 尤其是长上下文场景。原有 kernel 支持编译时固定的 k 值 (512/1024), 需要为每个 k 编译模块, 且性能有提升空间。该 PR 旨在提供运行时 k 的优化 kernel, 并作为 v1 的超集弃用旧版本。

实现拆解

1. 统一 JIT 模块加载: 修改 `topk.py`, 将 `_jit_topk_v2_module` 从参数化 (`topk` 编译时常量) 改为无参函数, 加载单一 `topk_v2` 模块。`plan_topk_v2` 和 `topk_transform_512_v2` 相应适配新接口, 且 transform 增加可选的 `out_raw_indices` 输出。
2. 核心 kernel 实现合并: 将旧版分散于 `register.cuh`、`streaming.cuh`、`cluster.cuh` 等多个头文件的实现, 重构为单一的 `topk_impl.cuh`。新实现采用 `TopKKernel` 类, 通过 `TopKProblem` 统一处理寄存器、流式、集群三种路径, 基于 plan 的 `cluster_threshold` 动态分派。
3. Plan-routed 分派和持久化集群池: 引入 `topk_plan` 内核生成元数据, 确定每个批次使用集群还是流式路径。持久化集群池 (30 个) 减少核函数发射, 适应小 batch 和大 batch 混合场景。
4. 测试和基准配套: 新增 `test_topk_v2.py` 覆盖所有分派路径的 172 个正确性测试; 新增 `bench_topk.py` 对 `jit_v1`、`jit_v2`、`flashinfer`、`torch` 进行 apples-to-apples 性能对比, 修正了 `flashinfer/torch` 基准缺少 page-table 变换的问题。

5. 删除旧实现：移除 `topk/register.cuh`、`topk/cluster.cuh`、`topk/streaming.cuh`、`topk/common.cuh`、`topk/ptx.cuh` 等旧版本头文件及相关引用。

关键文件：

- `test/registered/jit/deepseek_v4/test_topk_v2.py`（模块 测试覆盖；类别 test；类型 test-coverage；符号 `_assert_topk_close`, `_make_page_table`, `_invert`, `_reference`）：新增的正确性测试，系统覆盖所有 dispatch 模板（trivial, Register2, Register4, Streaming, Cluster）及其边界条件，共 172 个 case，是验证 kernel 正确性的核心文件。
- `test/registered/jit/benchmark/bench_topk.py`（模块 基准测试；类别 test；类型 test-coverage；符号 `_make_inputs`, `_make_p1_table`, `_build_fn`, `fn`）：新增的基准测试，对 jit_v1, jit_v2, flashinfer, torch 进行 apples-to-apples 性能比较，并修正了之前基准缺少 page-table 变换的问题。
- `python/sglang/jit_kernel/dsv4/topk.py`（模块 JIT Kernel 入口；类别 source；类型 core-logic；符号 `_jit_topk_v2_module`）：JIT kernel 的 Python 入口，简化了模块加载方式（去掉编译时 k），并适配 transform 和 plan 函数。
- `python/sglang/jit_kernel/include/sgl_kernel/deepseek_v4/topk_impl.cuh`（模块 核心 kernel；类别 other；类型 dependency-wiring）：核心 kernel 实现，将所有旧的注册、流式、集群路径合并为单一实现，采用运行时 top-k、MaxSmem 元编程等技术。
- `python/sglang/jit_kernel/csrc/deepseek_v4/topk_v2.cuh`（模块 Kernel 分派器；类别 other；类型 dependency-wiring）：Kernel 分派器和 host 端逻辑的修改，适配新实现并实现 plan-routed dispatch、persistent cluster pool 等优化。

关键符号：`_jit_topk_v2_module`, `plan_topk_v2`, `topk_transform_512_v2`, `_assert_topk_close`, `_reference`, `_run`, `_run_raw`, `benchmark`

关键源码片段

`test/registered/jit/deepseek_v4/test_topk_v2.py`

新增的正确性测试，系统覆盖所有 dispatch 模板（trivial, Register2, Register4, Streaming, Cluster）及其边界条件，共 172 个 case，是验证 kernel 正确性的核心文件。

```
# 最大允许错误数量（容忍 float16 直方图相等分值交换）
MAX_PERMIT_ERROR = 5
```

```
def _assert_topk_close(scores_cpu, ref_raw, our_raw, bs, seq_lens, k):
    """集合比较我们的 top-k 原始索引与 torch.topk 的，容忍相等分值交换。"""
    bad = 0
    for i in range(bs):
        L = int(seq_lens[i])
        ref, our = set(ref_raw[i]), set(our_raw[i])
        more, less = our - ref, ref - our
        if more or less:
            mv = sorted(scores_cpu[i, list(more)].tolist())
            lv = sorted(scores_cpu[i, list(less)].tolist())
            if mv != lv: # 不仅仅是相等交换 → 真正错误
                bad += len(more)
```

```

        print(
            f"b={i} L={L} k={k}: more={list(more)[:4]} less={list(less)[:4]} mv={mv[:3]} lv={lv[:3]}
            "
        )
        assert len(our) == min(k, L), f"b={i} L={L} k={k}: {len(our)} valid != {min(k, L)}"
        assert bad <= MAX_PERMIT_ERROR, f"{bad=} > {MAX_PERMIT_ERROR}"

```

测试用例配置覆盖所有模板和边界：

```

FIXED_CONFIGS = [
    # --- trivial (seq <= k) ---
    (8, 256), (16, 1024),
    # --- Register2 (level 0: max_seq <= 8192) ---
    (8, 4096), (8, 8192), (128, 8192), (300, 8192),
    # --- Register4 (level 1: 8192 < max_seq <= 16384) ---
    (8, 8193), (64, 16384), (256, 16384),
    # --- Streaming (level 2: max_seq > 16384, non-cluster) ---
    (8, 16385), (4, 32768), (16, 65535), (4, 65536), (100, 65536),
    # --- Cluster, fused small-batch kernel (batch <= 30) ---
    (1, 65537), (2, 131072), (8, 98304), (30, 131072),
    # --- Cluster, persistent pool + main kernel (30 < batch <= 128) ---
    (31, 131072), (40, 262144), (64, 196608), (128, 131072),
    # --- batch > 128 => non-cluster streaming ---
    (129, 131072), (200, 262144),
]

```

python/sglang/jit_kernel/dsv4/topk.py

JIT kernel 的 Python 入口，简化了模块加载方式（去掉编译时 k），并适配 transform 和 plan 函数。

```

@cache_once
def _jit_topk_v2_module():
    # v2 是通用的：topk (≤2048) 是运行时参数，不是编译时常量，
    # 因此单个模块即可服务所有 k 值。
    return load_jit(
        make_name("topk_v2"),
        cuda_files=["deepseek_v4/topk_v2.cuh"],
        cuda_wrappers=[
            ("topk_transform", "TopKKernel::transform"),
            ("topk_plan", "TopKKernel::plan"),
        ],
    )

# metadata 是 (batch+1, 2) int32: 第 0 行 = {cluster_threshold, num_cluster_items};
# 第 1..N 行为路由到持久化集群池项的 {batch_id, seq_len}。
_PLAN_METADATA_INTS_PER_BATCH = 2

def plan_topk_v2(seq_lens: torch.Tensor, static_threshold: int = 0) -> torch.Tensor:
    module = _jit_topk_v2_module()
    bs = seq_lens.shape[0]

```

```

metadata = seq_lens.new_empty(bs + 1, _PLAN_METADATA_INTS_PER_BATCH)
module.topk_plan(seq_lens, metadata, static_threshold)
return metadata

def topk_transform_512_v2(
    scores: torch.Tensor,
    seq_lens: torch.Tensor,
    page_tables: torch.Tensor,
    out_page_indices: torch.Tensor,
    page_size: int,
    metadata: torch.Tensor,
    out_raw_indices: Optional[torch.Tensor] = None,
) -> None:
    module = _jit_topk_v2_module()
    module.topk_transform(
        scores,
        seq_lens,
        page_tables,
        out_page_indices,
        page_size,
        metadata,
        out_raw_indices,
    )

```

评论区精华

Reviewer gemini-code-assist[bot] 指出了三个 PDL (Program Dependency Link) 触发问题 和一个类型不一致问题:

- topk_persistent_cluster_kernel 中 PDLTriggerSecondary 在写输出前调用, 导致主内核读取未初始化数据 (critical) 。
- topk_main_kernel 同样触发过早 (high) 。
- topk_small_batch_kernel 缺少 PDLTriggerSecondary, 可能引起死锁 (high) 。
- topk_plan 中 seq_lens 声明为 uint32_t* 但其余部分使用 int32_t (medium) 。

PR 作者未在讨论中回复, 但 PR 最终被合并, 相关问题未在后续提交中明确修复。

- PDL 触发时机可能导致竞态条件或死锁 (correctness): PR 作者未回应, PR 最终被合并, 这些问题实际影响未知。
- topk_plan 中 seq_lens 类型不一致 (correctness): 未解决, PR 已合并。

风险与影响

- 风险:
 - PDL 竞态条件风险: 若 review 指出的 PDL 触发时机问题真实存在, 可能导致输出错误或死锁, 影响 DeepSeek-V4 推理正确性。实际已合并, 风险可能被忽略或已在硬件 / 工具链层面保证安全。

- 类型不一致风险: seq_lens 使用 uint32_t 而非 int32_t, 在极端长度下可能产生符号转换问题, 但通常 seq_lens 非负, 实际影响较小。
- 硬件依赖风险: PDL 功能依赖 NVIDIA Blackwell 架构 (B200), 在非 PDL 硬件上可能回退到非 PDL 路径 (代码含条件编译), 但若未正确回退则在旧硬件上出错。
- 删除旧代码风险: 若某些场景仍依赖 v1 kernel (如特定 k 值或放弃 JIT 回退), v2 完全替代 v1 后需确保所有调用点已迁移。
- 影响:
 - 用户影响: 无 API 变化, 用户透明受益于 kernel 加速, 尤其长上下文解码场景。
 - 系统影响: DeepSeek-V4 推理的 top-k + page-table 变换在 B200 上带宽可达 3.0-3.6 TB/s (接近 HBM 峰值), 显著降低解码延迟。
 - 团队影响: 单一 kernel 实现替代分散的多个头文件, 降低维护成本, 但新 kernel 复杂性较高, 需要熟悉 CUDA JIT 和 PDL。
 - 风险标记: PDL trigger race condition, seq_lens type mismatch, hardware PDL dependency

关联脉络

- 暂无明显关联 PR