

PR #47463 完整报告

vllm-project/vllm

[Perf] Optimize `fused\_topk\_bias` for DSv4, 1.5~2x kernel performance improvement

合并时间: 2026-07-16 07:40

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/47463>

执行摘要

- 一句话: 优化 fused_topk_bias 数据类型转换, 性能提升 1.5-2 倍
- 推荐动作: 值得仔细阅读 CUDA kernel 的模板设计技巧——通过设备函数统一索引类型, 避免 Python 层拷贝。该模式可推广至其他需要异构类型输入的 kernel 优化。Python 侧逻辑简化也提升了可读性。

功能与动机

作为 DeepSeek V4 性能优化 (#45861) 的一部分, 原实现中当 hash_indices_table 非 None 时, Python 侧会调用 `.to()` 将 input_tokens 和 hash_indices_table 转为 indices_type (默认 int32), 造成不必要的显存拷贝与同步开销。将类型转换融合到 kernel 内可避免该冗余操作。

实现拆解

1. 简化 Python 侧逻辑 (vllm/model_executor/layers/fused_moe/router/fused_topk_bias_router.py): 移除以 indices_type 为基准的强制类型转换, 改为仅当 input_tokens 与 hash_indices_table 类型不一致时才做一次对齐转换, 同时移除相关注释。
2. 扩展 CUDA Kernel (csrc/libtorch_stable/moe/topk_softplus_sqrt_kernels.cu): 增加模板参数 HashIndType 以接受任意整数类型的 input_ids 和 tid2eid; 新增 load_index_as_int64 设备辅助函数, 将外部数据统一转为 int64_t 进行哈希索引查表与专家选择; 输出时强制转换为 IndType 以保持与原有接口一致。
3. 更新测试覆盖 (tests/kernels/moe/test_topk_softplus_sqrt.py): 将 hash_indices_table 和 input_ids 的数据类型从 torch.int32 改为 torch.long, 模拟真实 DSv4 场景下的 int64 输入, 确保 kernel 和 Python 侧都能正确处理。

关键文件:

- vllm/model_executor/layers/fused_moe/router/fused_topk_bias_router.py (模块 路由层; 类别 source; 类型 core-logic; 符号 fused_topk_bias): 移除冗余 dtype 转换, 简化控制流
- csrc/libtorch_stable/moe/topk_softplus_sqrt_kernels.cu (模块 CUDA 内核; 类别 source; 类型 core-logic; 符号 topkGatingSoftplusSqrt, load_index_as_int64): 新增模板参数, 内部统一用 int64 索引, 去除 Python 侧转换依赖

- tests/kernels/moe/test_topk_softplus_sqrt.py (模块 单元测试; 类别 test; 类型 test-coverage; 符号 test_fused_topk_softplus_sqrt_hash) : 测试输入类型从 int32 改为 int64, 匹配真实场景

关键符号: fused_topk_bias, topkGatingSoftplusSqrt, load_index_as_int64

关键源码片段

vllm/model_executor/layers/fused_moe/router/fused_topk_bias_router.py

移除冗余 dtype 转换, 简化控制流

```
def fused_topk_bias(
    hidden_states: torch.Tensor,
    gating_output: torch.Tensor,
    scoring_func: str,
    e_score_correction_bias: torch.Tensor,
    topk: int,
    renormalize: bool,
    indices_type: torch.dtype | None = None,
    input_tokens: torch.Tensor | None = None,
    hash_indices_table: torch.Tensor | None = None,
    routed_scaling_factor: float = 1.0,
):
    # 仅当 input_tokens 和 hash_indices_table 类型不一致时才做一次对齐转换
    # 原逻辑是对两者都向 indices_type 转, 现改为统一到 hash_indices_table 类型
    if (
        input_tokens is not None
        and hash_indices_table is not None
        and input_tokens.dtype != hash_indices_table.dtype
    ):
        input_tokens = input_tokens.to(dtype=hash_indices_table.dtype)

    # 其余路由逻辑保持不变, kernel 内部负责类型适配
    if not rocm_aiter_ops.is_fused_moe_enabled():
        ... # (softmax / sqrtsoftplus 分支 )
    else:
        ... # (fallback PyTorch 实现 )
```

csrc/libtorch_stable/moe/topk_softplus_sqrt_kernels.cu

新增模板参数, 内部统一用 int64 索引, 去除 Python 侧转换依赖

```
// 辅助函数: 将任意整数指针元素安全加载为 int64_t, 用于统一索引计算
template <typename HashIndType>
__device__ __forceinline__ int64_t load_index_as_int64(const HashIndType* ptr,
                                                         int64_t offset) {
    return static_cast<int64_t>(ptr[offset]);
}

// kernel 模板增加 HashIndType 参数
```

```

// topkGatingSoftplusSqrt<..., typename HashIndType>

// USE_HASH 分支：使用 hash 表预选 expert
if constexpr (USE_HASH) {
    // 通过 load_index_as_int64 统一将 input_ids 转为 int64_t
    const int64_t token_id = load_index_as_int64(input_ids, thread_row);
    const int64_t token_expert_offset = token_id * static_cast<int64_t>(k);
    // ... (softplus 计算相同)
    #pragma unroll
    for (int k_idx = 0; k_idx < k; ++k_idx) {
        const int expert = static_cast<int>(
            load_index_as_int64(tid2eid, token_expert_offset + k_idx));
        // 从 tid2eid 查询到的 expert 索引写入 indices (输出类型为 IndType)
        indices[idx] = static_cast<IndType>(expert);
        // ... (权重累加)
    }
}

```

评论区精华

Review 中 jeejeelee 提出问题：“能否在 routing kernel 中完成类型转换？”作者 yewentao256 采纳该建议，重写 CUDA kernel 使其内部处理类型转换，并更新 PR 描述中的性能数据。最终方案获得批准。

- 是否在 kernel 内完成类型转换 (design): 作者采纳建议，修改 CUDA kernel 使其内部自动处理类型转换并更新性能数据。

风险与影响

- 风险：
 1. kernel 模板膨胀：新增 HashIndType 模板参数可能增加编译实例化数量，但由于只有少数几种整数类型（int32、int64），影响可控。
 2. 索引性能：内部统一使用 int64 索引，对原本使用 int32 的设备可能引入微开销，但实测整体性能提升显著，表明该开销被消除的 Python 转换成本抵消。
 3. 输入类型假设：Python 侧去除了断言，若传入类型不一致且未命中新转换条件（如 indices_type 与 hash_indices_table.dtype 不同），kernel 可能收到混合类型，但当前逻辑仅依赖 hash_indices_table.dtype 转换，与原行为有细微差异，需确保调用方类型一致。
 4. 测试覆盖：测试仅使用 int64 输入，可能遗漏 int32 路径；但 kernel 已泛化支持，低风险。- 影响：对使用 DeepSeek V4 的推理任务，fused_topk_bias kernel 性能提升 1.5-2 倍，降低整体 MoE 路由延迟。对其他模型无影响（该 kernel 仅在启用 hash_indices_table 时被调用）。代码量净减少，维护成本降低。- 风险标记：kernel 模板复杂度增加，int64 索引潜在开销，测试仅覆盖 int64 路径

关联脉络

- PR #45861 Performance Optimization for Deepseek V4: 此 PR 是 DSv4 性能优化跟踪 Issue 中的子任务之一，直接关联该 Issue。