

PR #42943 完整报告

vllm-project/vllm

[CPU][RISC-V] Add VLEN=256 support to RVV attention kernels

合并时间: 2026-05-21 19:50

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/42943>

执行摘要

- 一句话: 将 RISC-V RVV attention 内核支持扩展到 VLEN=256, 性能提升达 3.58 倍
- 推荐动作: 值得精读, 特别是 RVVI() 宏的使用和 VLEN-agnostic 编程模式, 这是一个很好的 C++ 模板与宏结合的实践。对于 RISC-V 平台用户, 建议尽快部署以获取性能收益, 但需注意 VLEN>256 的兼容性风险。团队应评估并修复 Python 端检测过于宽松的问题。

功能与动机

根据 PR body, 目的是支持 VLEN=256 的 RISC-V 硬件 (如 Spacemit X100), 之前 VLEN=256 只能回退到标量 VEC 实现。扩展 RVV 内核支持可以释放 RISC-V 硬件的性能潜力。

实现拆解

1. 重构 RVV 内核 ([csrc/cpu/cpu_attn_rvv.hpp](#)): 将之前硬编码为 VLEN=128 的固定宽度 typedef (如 `fixed_vfloat32m2_t`) 和直接 intrinsics 调用替换为 VLEN-agnostic 的 RVVI() 宏和语义类型 (如 `fixed_fp32x8_t`)。条件编译门控从 `__riscv_v_min_vlen == 128` 改为 `__riscv_v_min_vlen == 128 || __riscv_v_min_vlen == 256`。
2. 更新类型定义 ([csrc/cpu/cpu_types_riscv_defs.hpp](#)): 新增 `fixed_u32x8_t` (LMUL_256 的 `uint32` 向量类型), 用于 VLEN=256 下的 BFloat16 fallback 路径中的位操作。
3. 修改条件编译生成器 ([csrc/cpu/generate_cpu_attn_dispatch.py](#)): 将 `include` 和 `dispatch` 宏中 RISC-V 分支从仅 VLEN==128 扩展为 VLEN==128 || VLEN==256, 合并了两个分支为单一 RVV+VEC+VEC16 块。
4. 调整 Python 分发 ([vllm/v1/attention/backends/cpu_attn.py](#)): 将 `_riscv_supports_rvv_vlen128` 重命名为 `_riscv_supports_rvv`, 并修改逻辑: 检查 `/proc/cpuinfo` 中是否存在 `zvl128b` 或 `zvl256b`, 同时确保不存在更大 VLEN (512/1024) 的标记, 以匹配 C++ 编译条件。

关键文件:

- `csrc/cpu/cpu_attn_rvv.hpp` (模块 RVV 内核; 类别 source; 类型 core-logic; 符号 `load_row8_B_as_f32`, `gemm_micro_rvv_fma_Mx8_Ku4`): 核心注意力内核, 从硬编码 VLEN=128 重构为 VLEN-agnostic, 使用 RVVI() 宏和语义类型定义。
- `vllm/v1/attention/backends/cpu_attn.py` (模块 注意力后端; 类别 source; 类型 core-logic; 符号 `_riscv_supports_rvv_vlen128`, `_riscv_supports_rvv`, `_get_attn_isa`): Python 分发逻辑, 检测 RVV 支持并选择 ISA, 函数 `_riscv_supports_rvv` 被重命名并调整

逻辑。

- `csrc/cpu/generate_cpu_attn_dispatch.py` (模块 分发生成器; 类别 `source`; 类型 `core-logic`; 符号 `generate_header_file`) : 条件编译生成器, 合并了 RVV dispatch 宏的逻辑, 从仅 `VLEN=128` 扩展到 `128` 或 `256`。
- `csrc/cpu/cpu_types_riscv_defs.hpp` (模块 类型定义; 类别 `source`; 类型 `core-logic`; 符号 `fixed_u32x8_t`) : 添加了 `uint32` 向量类型 `fixed_u32x8_t`, 用于 `VLEN=256` 下的特定操作。

关键符号: `_riscv_supports_rvv`, `load_row8_B_as_f32`, `gemm_micro_rvv_fma_Mx8_Ku4`, `generate_header_file`

关键源码片段

`csrc/cpu/cpu_attn_rvv.hpp`

核心注意力内核, 从硬编码 `VLEN=128` 重构为 `VLEN-agnostic`, 使用 `RVVI()` 宏和语义类型定义。

```
// cpu_attn_rvv.hpp: VLEN-agnostic 的注意力内核
//
// 使用 RVVI() 宏使 intrinsics 调用独立于 VLEN,
// semantic type (fixed_fp32x8_t) 自动映射到正确 LMUL

#if defined(__riscv_v_min_vlen) && \
    (__riscv_v_min_vlen == 128 || __riscv_v_min_vlen == 256)

#include "cpu_attn_impl.hpp"
#include "cpu_types_riscv_defs.hpp"
#include <riscv_vector.h>

namespace cpu_attention {
namespace {

// 加载 8 个行元素为 FP32 向量
// 返回 fixed_fp32x8_t, VLEN=128 时为 m2 (8 个元素), VLEN=256 时为 m1 (8 个元素)
template <typename kv_cache_t>
FORCE_INLINE fixed_fp32x8_t load_row8_B_as_f32(const kv_cache_t* p);

// float 特化: 直接加载
template <>
FORCE_INLINE fixed_fp32x8_t load_row8_B_as_f32<float>(const float* p) {
    // RVVI(__riscv_vle32_v_f32, LMUL_256) 在 VLEN=128 时展开为 vle32_v_f32m2,
    // VLEN=256 时展开为 vle32_v_f32m1
    return RVVI(__riscv_vle32_v_f32, LMUL_256)(p, 8);
}

// Half 特化: 加载 8 个 float16 并扩宽到 float32
template <>
FORCE_INLINE fixed_fp32x8_t load_row8_B_as_f32<c10::Half>(const c10::Half* p) {
```

```

#ifdef __riscv_zvfh
// 加载 8 个 float16 (LMUL_128 保证 8 个元素)
fixed_fp16x8_t h = RVVI(__riscv_vle16_v_f16, LMUL_128)(
    reinterpret_cast<const _Float16*>(p), 8);
// 扩宽为 float32 (LMUL_256 自动处理 LMUL)
return RVVI(__riscv_vfwcvt_f_f_v_f32, LMUL_256)(h, 8);
#else
// 无 Zvfh 硬件时的标量 fallback
alignas(16) float tmp[8];
for (int i = 0; i < 8; ++i) {
    tmp[i] = static_cast<float>(p[i]);
}
return RVVI(__riscv_vle32_v_f32, LMUL_256)(tmp, 8);
#endif
}

```

vllm/v1/attention/backends/cpu_attn.py

Python 分发逻辑，检测 RVV 支持并选择 ISA，函数 `_riscv_supports_rvv` 被重命名并调整逻辑。

cpu_attn.py: Python 端 RVV 支持检测

```

@functools.lru_cache(maxsize=1)
def _riscv_supports_rvv() -> bool:
    """检测 C++ RVV attention 内核是否可用。

```

内核现在支持 VLEN=128 和 256。CMake 根据 `/proc/cpuinfo` 中最大的 `zvl<N>b` 设置 `__riscv_v_min_vlen`，因此检查是否存在 `zvl128b` 或 `zvl256b`，并确保没有更大 VLEN (512/1024) 的广告以避免误判。

```

"""

```

```

try:
    with open("/proc/cpuinfo") as f:
        cpuinfo = f.read()
except OSError:
    return False
# 支持 128 或 256，但排除 512/1024 (它们会向前兼容广告小 VLEN)
return any(f"zvl{n}b" in cpuinfo for n in (128, 256)) and all(
    f"zvl{n}b" not in cpuinfo for n in (512, 1024)
)

```

在 ISA 选择中使用 (简化) :

```

def _get_attn_isa(...):
    ...
    if supports_riscv and _riscv_supports_rvv():
        return "rvv"
    ...

```

评论区精华

gemini-code-assist[bot] 在 review 中提出一个高优先级问题：Python 端的 `__riscv_supports_rvv` 在 `VLEN>256` 的系统（如 `VLEN=512`）上可能误判。因为 `/proc/cpuinfo` 通常向前兼容地包含所有较低 `VLEN` 的标记（如 `zvl128b` 和 `zvl256b`），导致 Python 端认为 RVV 内核可用，但 C++ 端因 `__riscv_v_min_vlen==512` 而不会编译 RVV 内核，最终在首次 `attention` 调用时触发 `TORCH_CHECK` 失败。该评论未得到回复，但 PR 仍被批准合并，风险未在代码层面解决。

- Python 端 `VLEN` 检测可能过于宽松 (correctness): 未回复，PR 被批准合并，风险未解决。

风险与影响

- 风险：条件编译不匹配风险：Python 端检测 `/proc/cpuinfo` 中 `zvl128b` / `zvl256b` 的存在，但 C++ 内核依赖于编译时 `__riscv_v_min_vlen` 的精确值。在 `VLEN=512` 的系统上，Python 端会错误地返回 `True`（因为 `zvl128b` 和 `zvl256b` 都存在），而 C++ 端不会编译 RVV 内核，导致运行时崩溃。此外，虽然添加了 `uint32` 类型，但未添加其他 `LMUL` 变体，可能限制未来对其他 `VLEN` 的扩展。测试仅在 `VLEN=256` 设备上进行，未覆盖 `VLEN=512` 的回归场景。
- 影响：
 - 用户影响：RISC-V `VLEN=256` 用户（如 Spacemit X100）获得显著性能提升（`prefill` 3.01x, `mixed` 3.58x）；`VLEN=128` 用户无变化；`VLEN>256` 用户可能因上述风险遇到运行时错误。
 - 系统影响：CPU 注意力后端分发路径增加了一个 `ISA` 分支，但逻辑清晰，维护成本低。统一宏使未来支持更多 `VLEN` 更容易。
 - 团队影响：代码复杂度略有降低，但需要额外关注条件编译的一致性。
 - 风险标记：Python 端检测过于宽松，条件编译不匹配，`VLEN>256` 系统潜在回归

关联脉络

- 暂无明显关联 PR