

# PR #50387 完整报告

vllm-project/vllm

[CPU] Bump up CPU kernels to latest version

合并时间: 2026-07-30 19:18

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/50387>

## PR 分析报告: 同步 SGLang 上游 CPU 内核

### 执行摘要

本次 PR 将 vLLM 的 CPU 推理内核同步至 SGLang 上游最新版本, 合并了 vLLM 特有的 MXFP4 MoE、AMX GDN 等补丁, 并拒绝了一个上游回归。核心优化是 GDN 预填充路径采用池化状态索引直接读写, 避免一次额外拷贝, 带来 3%-8% 的性能提升。然而, review 中指出的两个低严重性 bug (bias 应用错误和空指针算术 UB) 未被修复, 可能埋下隐患。

### 功能与动机

自上次同步 (#41924) 以来, vLLM 团队在 CPU 内核上积累了 MXFP4 权重量化的 MoE 支持、AMX GDN 调度、ISA 可移植 BLAS fallback、RISC-V 标量 / 向量支持、chunked-prefill 的 `has_initial_state` 修复以及 DFlash 推测解码等独家补丁。保持与上游一致有助于降低维护成本并引入上游的优化 (如快速 silu 近似、池化状态索引)。此外, 在合并过程中发现上游的 MXFP4scale 有效性检查 `TORCH_CHECK` 条件被错误放宽, 需拒绝此回归以确保量化正确性。

### 实现拆解

1. 同步上游代码并融合 vLLM 补丁: 整体替换 `csrc/cpu/sgl-kernels/` 目录, 将每个 vLLM 特有修改重新应用到新代码中, 解决冲突。这部分涉及 `moe.cpp`、`moe_int8.cpp`、`moe_fp8.cpp`、`conv.cpp`、`fla.cpp` 等多个文件的行级调整。
2. GDN 预填充优化: 在 `fla.cpp` 中, 采纳 SGLang 的新 `chunk_gated_delta_rule_cpu` 签名, 要求传入 `initial_state_indices`, 使 CPU 路径直接读写池化的 `ssm_state` 缓冲区, 而不再 gather 出本地数组后写回。这消除了每次预填充的拷贝开销, 并配合函数内部的其他融合 (`kkt_solve` 与 `recompute_w_u` 融合, `recompute_w_u` 与 `update_v` 融合) 进一步降低延迟。
3. 拒绝上游回归 + 补全回退路径: 发现并拒绝了一个在上游版本中无意削弱的 `TORCH_CHECK` 条件 (MXFP4 scale 检查), 还原为原严格版本。同时, 为 ARM 和非 AVX512 架构编写了基于 `at::vec::Vectorized` 的向量化 fallback, 替换之前更原始的标量实现, 并修复了其中 `kNumHead/kHeadDim` 条件判断错误。
4. MoE 内核能力增强: 在 `moe.cpp` 和 `moe.h` 中, `fused_experts_kernel_impl` 新增 `alpha`、`limit`、`act_func` (支持 `silu_and_mul`、`swiglu`、`gelu_and_mul`) 和 `with_bias` 参数, 使 MoE 层能处理不同的门控激活函数和偏置注入。同时将 AVX512 路径的 `silu` 替换为 `_mm512_rcp14_silu_ps` 快速近似, 减少延迟。

5. 测试与构建适配：更新 torch\_bindings.cpp 中的算子注册匹配新签名；修改 tests/cpu/gdn/ops/test\_cpu\_gdn\_ops.py 以传入 initial\_state\_indices 参数，并限制仅支持 head\_dim 为 64 或 128 的情形（与上游一致）。

### csrc/cpu/sgl-kernels/fla.cpp

GDN 内核完全重写，引入池化状态索引、循环融合、AMX BF16 状态更新等关键优化，是本次 PR 的核心性能改进所在。

```
// fla.cpp: pack_vnni2 — 将 FP32 数据转换为 VNNI 格式 (BF16) 并缩放
// 这是 AMX 优化的关键步骤：将布局从 [K/2, 2, N] 转换为 [K/2, N, 2] 以匹配 AMX 输入格式
// 同时将 src 乘以 exp(g_last) 实现门控衰减，避免后续单独处理
template <typename scalar_t, int K, int N>
void pack_vnni2(scalar_t* __restrict__ dst, float* __restrict__ src,
               const float g_last, int ld_src, int ld_dst) {
    static_assert(K % 32 == 0);
    static_assert(N % 32 == 0);

    const float scale = std::exp(g_last);
#ifdef CPU_CAPABILITY_AVX512
    constexpr int KB = K / 2;
    constexpr int NB = N / 32;
    __m512i s0, s1, d0, d1;
    __m512 vd = _mm512_set1_ps(scale);

    // Unroll 遍历 KB * NB 个 tile，每个 tile 处理 32 个元素
    const auto trans = [&](auto i) {
        constexpr int kb = i / NB;
        constexpr int nb = i % NB;

        // 从 [K/2, 2, N/32, 32] 布局读取
        constexpr int k0 = kb * 2 + 0;
        constexpr int k1 = kb * 2 + 1;
        __m512 v00 = _mm512_loadu_ps(src + k0 * ld_src + nb * 32);
        __m512 v01 = _mm512_loadu_ps(src + k0 * ld_src + nb * 32 + 16);
        __m512 v10 = _mm512_loadu_ps(src + k1 * ld_src + nb * 32);
        __m512 v11 = _mm512_loadu_ps(src + k1 * ld_src + nb * 32 + 16);
        s0 = (__m512i)_mm512_cvtne2ps_pbh(v01, v00); // 两路 BF16 压缩
        s1 = (__m512i)_mm512_cvtne2ps_pbh(v11, v10);

        std::tie(d0, d1) = transpose_2x32_16bit(s0, s1); // 转置为 [N/32, 32, 2]
        _mm512_storeu_si512(dst + kb * ld_dst * 2 + nb * 32 * 2, d0);
        _mm512_storeu_si512(dst + kb * ld_dst * 2 + nb * 32 * 2 + 32, d1);

        // 原地缩放 src 以避免后续再遍历
        _mm512_storeu_ps(src + k0 * ld_src + nb * 32, _mm512_mul_ps(v00, vd));
        _mm512_storeu_ps(src + k0 * ld_src + nb * 32 + 16, _mm512_mul_ps(v01, vd));
        _mm512_storeu_ps(src + k1 * ld_src + nb * 32, _mm512_mul_ps(v10, vd));
        _mm512_storeu_ps(src + k1 * ld_src + nb * 32 + 16, _mm512_mul_ps(v11, vd));
    };
};
```

```

    Unroll<KB * NB>{}(trans);
#else
    // 非 AVX512 回退：简单循环，等价但未微架构优化
    for (int k = 0; k < K; k += 2) {
        for (int n = 0; n < N; ++n) {
            const float v0 = src[(k + 0) * ld_src + n];
            const float v1 = src[(k + 1) * ld_src + n];
            dst[(k >> 1) * ld_dst * 2 + n * 2 + 0] = static_cast<scalar_t>(v0);
            dst[(k >> 1) * ld_dst * 2 + n * 2 + 1] = static_cast<scalar_t>(v1);
            src[(k + 0) * ld_src + n] *= scale;
            src[(k + 1) * ld_src + n] *= scale;
        }
    }
#endif
}

```

### csrc/cpu/sgl-kernels/moe.cpp

MoE 核心内核实现，新增 bias 支持和多种激活函数，优化 tinygemm 的 silu 融合，是 MoE 前向的正确性和性能关键。

```

// moe.cpp: tinygemm_kernel_nn2 的 storec 改造 — 使用快速 rcp14 silu 近似
// 原实现: x0 = x0 / (one + x0.neg().exp_u20())
// 新实现: 使用 _mm512_rcp14_silu_ps 指令，降低延迟且数值足够接近
// 同时输出格式直接从两个 FP32 向量打包为 BF16 存储
auto storec = [&](auto i) {
    constexpr int row = i / COLS;
    constexpr int col = i % COLS;
    if constexpr (col % 2 == 0) { // 每两列合并为一个 AVX512 存储
        __m512 x0 = vc0[row * COLS + col + 0];
        __m512 x1 = vc0[row * COLS + col + 1];
        __m512 y0 = vc1[row * COLS + col + 0];
        __m512 y1 = vc1[row * COLS + col + 1];
        // 快速 silu: silu(x) = x * sigmoid(x) 的倒数近似
        x0 = _mm512_mul_ps(_mm512_rcp14_silu_ps(x0), y0);
        x1 = _mm512_mul_ps(_mm512_rcp14_silu_ps(x1), y1);
        // 直接打包为 BF16 存储，去掉中间转换步骤
        _mm512_storeu_si512(
            reinterpret_cast<__m512i*>((C + row * ldc + col * 16)),
            (__m512i)(_mm512_cvtne2ps_pbh(__m512(x1), __m512(x0))));
    }
};

```

### 评论区精华

- Bias 应用错误 (depthfirst-app[bot])：当不使用 brgemm 且激活为 silu\_and\_mul 时，融合 tinygemm 路径输出写入 ic1，但后续 bias 添加针对 C0/C1，导致 bias 静默丢失。严重性：中高，未回复未修复。
- 空指针算术 UB (depthfirst-app[bot])：即使 with\_bias=false，代码仍计算 w1\_bias + expert\_id \* 2 \* N (w1\_bias 为 nullptr)，属未定义行为。严重性：中高，未回复未修复。

- 此外，CI 流程多次触发（4 次 /ci run），最终 Buildkite 构建通过，测试结果覆盖 12 个量化模型，大多数字节一致。

## 风险与影响

- 风险：两个未解决的 review 问题可能导致特定场景下的推理错误或编译器优化破坏保护逻辑；2/12 测试模型的输出偏移暗示量化内核数值精度有变化，可能影响生产的一致性。
- 影响：所有 CPU 推理用户都需要重新构建，并建议在目标任务上重新验证模型精度。GDN 用户可立即获得 3%-8% 的性能收益，但需要额外注意是否触发带 bias 的 MoE 门控路径（若触发则会得到错误结果）。
- 后续：作者应在后续 PR 中修复 bias 和 UB 问题，并扩展测试覆盖非量化路径和 RISC-V 架构。

## 关联脉络

- #41924：本次同步的上次基准。所有 vLLM 特有补丁均在此之后累积。
  - #47301（GDN streaming derenderer 前端）：虽然不直接相关，但 GDN 预填充优化可能与此类流式推理场景协同产生更大收益。
  - 整体上，这个 PR 标志着 vLLM CPU 内核与 SGLang 上游的持续对齐策略，未来每次上游更新都需经过类似的融合流程。