

PR #46202 完整报告

vllm-project/vllm

[CPU] Enable chunked prefill and prefix caching for qwen3.5

合并时间: 2026-06-25 11:49

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/46202>

执行摘要

- 一句话: CPU 线性注意力模型支持 chunked prefill 和前缀缓存
- 推荐动作:
 - 值得精读 `_zero_block_ids` 的实现, 它参考 GPU `KVBlockZeroer` 进行适配, 是混合注意力模型正确性的关键。
 - `batch_memcpy` 回退展示了如何在缺乏 Triton 时用 `ctypes` 兼容核心特性, 值得其他后端参考。
 - `conv.cpp` 的 `has_initial_state` 修复是 chunked prefill 正确的核心。
 - 建议关注 `depthfirst-app[bot]` 提出的空指针问题是否在后续修复。

功能与动机

之前 CPU 平台上对线性注意力模型无条件下禁用 chunked prefill 和 prefix caching (见 `vllm/platforms/cpu.py`), 因为担心正确性问题。经过验证底层内核修复 (`conv` 忽略 `has_initial_state`) 和模型运行器调整, 这些功能现在可以安全启用, 从而提升 CPU 推理吞吐并降低首 token 延迟。PR body 指出目的为移除过度保守的 guard 并修复混合模型中部分写入块的问题。

实现拆解

1. 移除 `vllm/platforms/cpu.py` 中的禁用 Guard 在 `CpuPlatform.check_and_update_config` 中删除了根据 AMX 支持和线性注意力层数量强制设置 `enable_prefix_caching=False` 和 `enable_chunked_prefill=False` 的代码块。
2. 修复 `CPUModelRunner._zero_block_ids` 从空操作改为对 `FullAttentionSpec` 类型的 KV 缓存块数据清零, 避免混合注意力模型中非全注意力块的部分写入导致脏数据影响计算。使用 `data_ptr` 去重防止重复清零。
3. 修复 C++ `conv` 内核忽略 `has_initial_state` 在 `csrc/cpu/sgl-kernels/conv.cpp` 的 `causal_conv1d_fwd_varlen_kernel_impl` 中, 将硬编码的 `nullptr` 和 `false` 替换为从参数传递的 `conv_states` 和 `has_initial_state[bs]`, 使 `varlen` 卷积在继续 prefill chunk 时能正确使用前一步的卷积状态。
4. 添加 `batch_memcpy_kernel` 的 CPU Fallback 在 `vllm/utils/cpu_triton_utils.py` 中实现 `_batch_memcpy_impl`, 基于 `ctypes.memmove` 批量拷贝内存, 用于 `triton-cpu` 不可用时替代 `mamba_utils.batch_memcpy_kernel`。在 `cpu_model_runner._postprocess_triton`

中注册该回退。

5. 添加回归测试 - 在 `tests/kernels/mamba/cpu/test_cpu_gdn_ops.py` 中新增 `test_chunk_gated_delta_rule_cpu_two_call_split` 等测试，模拟两次调度步骤间的状态传递，验证内核跨 chunk 的正确性。 - 在 `tests/v1/e2e/test_cpu_linear_attn_chunked_prefix.py` 中新增端到端测试，使用 `check_logprobs_close` 比较 chunked prefill 与 full prefill 的 logprobs，以及前缀缓存命中与冷缓存的输出。 - CI 配置调整：增大 CPU 语言 / 池化测试超时，将线性注意力测试移动至 triton-cpu job。

关键文件：

- `vllm/platforms/cpu.py` (模块 平台层；类别 source；类型 core-logic；符号 `check_and_update_config`)：移除强制禁用 chunked prefill 和 prefix caching 的 guard，这是功能启用的核心开关。
- `vllm/v1/worker/cpu_model_runner.py` (模块 推理引擎；类别 source；类型 core-logic；符号 `_zero_block_ids`)：修复 `_zero_block_ids` 对混合注意力模型中非全注意力块的部分写入问题，是正确性关键。
- `vllm/utils/cpu_triton_utils.py` (模块 工具库；类别 source；类型 core-logic；符号 `_batch_memcpy_impl, batch_memcpy_kernel`)：新增 `batch_memcpy` CPU 回退实现，解决 triton-cpu 不可用时的功能缺失。
- `csrc/cpu/sgl-kernels/conv.cpp` (模块 CPU 内核；类别 source；类型 core-logic；符号 `LAUNCH_TINYGEMM_VARLEN_KERNEL, causal_conv1d_fwd_varlen_kernel_impl`)：修复 `varlen conv` 内核忽略 `has_initial_state`，使 chunked prefill 时卷积状态正确传递。
- `tests/kernels/mamba/cpu/test_cpu_gdn_ops.py` (模块 GDN 算子；类别 test；类型 test-coverage；符号 `test_chunk_gated_delta_rule_cpu_two_call_split, test_causal_conv1d_torch_two_call_split, test_causal_conv1d_fwd_cpu_two_call_split, test_batch_memcpy_cpu_fallback`)：新增两步拆分测试，验证 GDN 内核跨 chunk 状态传递正确性，为 chunked prefill 提供 kernel 级回归。
- `tests/v1/e2e/test_cpu_linear_attn_chunked_prefix.py` (模块 线性注意力；类别 test；类型 test-coverage；符号 `test_chunked_prefill_matches_full_prefill, test_prefix_cache_hit_matches_cold_cache, full_prefill_refs`)：新增 chunked prefill 与 prefix caching 端到端回归测试，确保精度和缓存命中正确。

关键符号：`check_and_update_config, _zero_block_ids, _batch_memcpy_impl, batch_memcpy_kernel, causal_conv1d_fwd_varlen_kernel_impl, test_chunk_gated_delta_rule_cpu_two_call_split, test_causal_conv1d_torch_two_call_split, test_causal_conv1d_fwd_cpu_two_call_split, test_batch_memcpy_cpu_fallback, test_chunked_prefill_matches_full_prefill, test_prefix_cache_hit_matches_cold_cache`

关键源码片段

`vllm/platforms/cpu.py`

移除强制禁用 chunked prefill 和 prefix caching 的 guard，这是功能启用的核心开关。

```
# vllm/platforms/cpu.py (check_and_update_config excerpt)
```

```

# 下面的 if 块已在 PR#46202 中删除:
# if torch.cpu._is_amx_tile_supported() and (
# model_config is not None
# and model_config.get_num_layers_by_block_type(
# parallel_config, "linear_attention")
# ) > 0
# ):
# cache_config.enable_prefix_caching = False
# scheduler_config.enable_chunked_prefill = False
# logger.warning_once(
# "Disabled unsupported prefix caching and chunked prefill \"
# \"for linear attention on AMX CPU platforms.\"
# )
# 现在这两个配置项可以根据用户设置正常使用。

```

vllm/v1/worker/cpu_model_runner.py

修复 `_zero_block_ids` 对混合注意力模型中非全注意力块的部分写入问题，是正确性关键。

```

def _zero_block_ids(self, block_ids: list[int]) -> None:
    # Zero full-attention blocks to prevent stale data corruption on partial writes.
    # `FullAttentionSpec` 过滤排除了 encoder-only 层，避免误清零
    seen_ptrs: set[int] = set()
    for group in self.kv_cache_config.kv_cache_groups:
        # 只处理 FullAttentionSpec 的 KV 缓存组
        if not isinstance(group.kv_cache_spec, FullAttentionSpec):
            continue
        for layer_name in group.layer_names:
            ctx = self.compilation_config.static_forward_context.get(layer_name)
            if ctx is None:
                continue
            kv = ctx.kv_cache
            if not isinstance(kv, torch.Tensor):
                continue
            # 避免重复清零同一 Tensor
            if kv.data_ptr() in seen_ptrs:
                continue
            seen_ptrs.add(kv.data_ptr())
            for block_id in block_ids:
                kv[block_id].zero_() # 清零块数据

```

评论区精华

- 正确性验证: Reviewer @bigPYJ1151 提出“Did you check the accuracy?”，作者 @tianmu-li 回复已验证，并在调节 `max_num_batched_tokens=128` 后发现精度问题后修复，最终确认 `chunked prefill` 与 `full prefill` 的 `logprobs` 接近（`gsm8k` 准确率 0.92）。
- 空指针风险: 自动化检查工具 `depthfirst-app[bot]` 在 `csrc/cpu/sgl-kernels/conv.cpp` 指出：新代码仅在 `has_conv_states` 真时访问 `has_initial_state[bs]`，但 `has_initial_state` 可能是空指针，建议增加 `has_initial_state != nullptr` 检查。该建议尚未看到对应修复提交。

- 正确性验证 (correctness): 解决, 准确率 0.92, 与 full prefill 一致。
- has_initial_state 空指针风险 (correctness): 未在 PR 中修复, 可能需要后续 PR 处理。

风险与影响

- 风险:
 - 空指针风险: 在 `csrc/cpu/sgl-kernels/conv.cpp` 的 `causal_conv1d_fwd_varlen_kernel_impl` 中, 新代码 `has_initial_states_value = has_conv_states ? has_initial_state[bs] : false` 仅检查 `has_conv_states` 而不检查 `has_initial_state` 是否为空指针; 如果 `conv_states` 非空但 `has_initial_state` 为空, 则导致解引用空指针。当前调用点可能总是同时传参, 但缺乏防御性检查。
 - 部分写入脏数据: `_zero_block_ids` 的修复解决了混合注意力模型在 chunked prefill 时的脏数据问题, 但仅对 FullAttentionSpec 类型清零; 若未来有其他注意力规格, 可能需要更新。
 - 性能回退: `batch_memcpy` 使用 Python 循环加 `ctypes.memmove` 实现, 性能远低于原生 Triton 内核, 但在 triton-cpu 不可用时 fallback 路径触发, 影响前缀缓存命中的 Mamba 状态拷贝, 频率不高。
 - 测试覆盖有限: 端到端测试仅覆盖 Qwen3.5-0.8B 单模型和有限 prompt 长度, 未涵盖所有线性注意力变体或更大模型。
- 影响:
 - 用户: 在 CPU 上运行 Qwen3.5 等线性注意力模型时, 可启用 `--enable-chunked-prefill` 和 `--enable-prefix-caching`, 显著提升吞吐并减少首 token 延迟 (尤其长上下文场景)。
 - 系统: 移除全局面上的限制后, 所有 CPU 平台 (无论是否 AMX) 均尝试启用这些功能; AMX 平台之前被禁用, 现在启用, 可能增加 AMX 单元负载。
 - 团队: 涉及 CPU 后端的多个模块 (平台配置、模型运行器、工具函数、C++ 内核、测试), 需确保后续线性注意力模型不回归。测试用例为后续开发提供了回归保障。
 - 风险标记: 空指针风险, 性能回退, 测试覆盖有限

关联脉络

- PR #40172 `batch_memcpy_kernel` in `mamba_utils`: 本 PR 添加的 `batch_memcpy` CPU fallback 是对 upstream PR #40172 的补充, 该 PR 在 `mamba_utils` 中引入了 `batch_memcpy_kernel` 但未提供 CPU 回退。
- PR #41025 Accuracy benchmark for models: fadara01 在评论中建议使用该 PR 的 accuracy benchmark 验证正确性。