

PR #41217 完整报告

vllm-project/vllm

[ROCm][Deepseek] dsv3.2 further optimization

合并时间: 2026-05-01 22:06

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/41217>

执行摘要

- 一句话: ROCm DeepSeek V3.2 MLA 注意力后端深度优化
- 推荐动作: 值得精读: 本 PR 展示了 ROCm 平台上高效 MLA 注意力 kernel 的设计模式, 包括 tile 策略、metadata 提前构建、融合 kernel 替换等。建议重点阅读 `rocm_aiter_mla_sparse.py` 中的 `_convert_req_index_to_global_index_kernel` 和 `_forward_mla` 实现, 以及 `deepseek_v2.py` 中的 ROCm 分支设计。

功能与动机

参考 PR body 及替换的 #32649, 本 PR 旨在进一步提升 ROCm 平台上 DeepSeek V3.2 模型的推理性能, 通过 kernel 级优化和架构调整降低 MLA 注意力机制的延迟。GSM8K 测试分数从 baseline 提升至 0.9477。

实现拆解

1. 新增 Triton 索引 kernel: 在 `rocm_aiter_mla_sparse.py` 中新增 `_convert_req_index_to_global_index_kernel`, 替换原从 `flashmla_sparse` 导入的 kernel, 采用 tile 策略优化缓存行, 实现更高效的请求索引到全局索引的转换。
2. Ragged metadata 提前构建: 新增 `generate_sparse_seqlen_kernel` 等 Triton kernel, 在 `metadata_builder` 阶段即生成稀疏序列长度信息, 避免在 `decoder forward` 中重复计算。
3. FP8 稀疏 MLA decode 支持: 在 `rocm_aiter_mla_sparse.py` 中实现 `decode` 阶段对 FP8 KV cache 和 logits 的稀疏注意力计算, 并通过 `current_platform.is_rocm()` 条件分支启用。
4. ROCm 专用 RoPE 路径: 在 `deepseek_v2.py` 的 `SparseAttnIndexer.forward` 中, 为 ROCm 分流旋转位置编码逻辑, 简化张量操作 (直接使用 `rotary_emb` 原地改写), 避免不必要的 `split` 和 `cat`。
5. 替换自定义 ops 调用: 将 `indexer_k_quant_and_cache` 从 `_custom_ops` 调用改为 `indexer_k_quant_and_cache_triton` (Triton 实现), 并支持 FP8 dtype 视图转换; 移除对 `is_cuda_alike` 的强依赖, 清理基础设施。
6. 融合 paged MQA logits: 将原两阶段 `stage1 + sum` 的 fp8 paged MQA logits 实现替换为 `deepgemm_fp8_paged_mqa_logits` 单 kernel 调用, 减少显存中间读写。
7. 扩展 kernel block 大小: 在 `indexer.py` 中将 ROCm 平台支持的 kernel block sizes 改为 `[1, 64]`, 使 `indexer` 能兼容 block size 1 的稀疏索引。

8. 文档更新：在 `attention_backends.md` 中更新 `ROCM_AITER_MLA_SPARSE` 的 `block sizes` 和 `FP8` 支持信息。

关键文件：

- `vllm/v1/attention/backends/mla/rocm_aiter_mla_sparse.py`（模块 `MLA` 后端；类别 `source`；类型 `core-logic`；符号 `_convert_req_index_to_global_index_kernel`, `triton_convert_req_index_to_global_index`, `generate_sparse_seqlen_kernel`, `generate_sparse_seqlen_triton`）：核心修改文件：新增 Triton 索引 `kernel`、ragged metadata 构建、FP8 sparse `MLA` 支持，实现了主要性能优化。
- `vllm/model_executor/models/deepseek_v2.py`（模块 `模型执行`；类别 `source`；类型 `data-contract`）：修改 `Indexer forward` 方法，为 `ROCm` 创建专用 `RoPE` 路径，避免影响其他平台。
- `vllm/v1/attention/ops/rocm_aiter_mla_sparse.py`（模块 `稀疏操作`；类别 `infra`；类型 `infrastructure`）：基础设施调整：移除对 `CUDA ops` 的依赖，使用 Triton 实现 `indexer cache` 操作，并集成融合 `kernel` 调用。
- `vllm/v1/attention/backends/mla/rocm_aiter_mla.py`（模块 `MLA 注意力`；类别 `source`；类型 `core-logic`）：添加 `_AITER_UNSUPPORTED_HEADS` 列表，在 `get_mla_padded_q` 中增加断言以防止 `head=32` 时非法访问。
- `vllm/v1/attention/backends/mla/indexer.py`（模块 `索引器`；类别 `source`；类型 `core-logic`）：修改 `get_supported_kernel_block_sizes` 返回 `[1,64]` 以支持 `block size 1`，适配 `ROCm` 稀疏 `MLA`。
- `docs/design/attention_backends.md`（模块 `文档`；类别 `docs`；类型 `documentation`）：更新文档反映 `ROCM_AITER_MLA_SPARSE` 新增的 `block sizes` 和 `FP8` 支持。

关键符号： `_convert_req_index_to_global_index_kernel`,
`triton_convert_req_index_to_global_index`, `generate_sparse_seqlen_kernel`,
`generate_sparse_seqlen_triton`, `_forward_bf16_kv`, `_forward_mla`,
`SparseAttnIndexer.forward`, `AiterMLAHelper.get_mla_padded_q`,
`DeepseekV32IndexerBackend.get_supported_kernel_block_sizes`

关键源码片段

`vllm/v1/attention/backends/mla/rocm_aiter_mla_sparse.py`

核心修改文件：新增 Triton 索引 `kernel`、ragged metadata 构建、FP8 sparse `MLA` 支持，实现了主要性能优化。

```
# vllm/v1/attention/backends/mla/rocm_aiter_mla_sparse.py
# 核心 kernel: 将 token_indices (局部索引) 通过 block_table 映射为全局 paged_kv_indices
@triton.jit
def _convert_req_index_to_global_index_kernel(
    req_id_ptr, # int32 [num_tokens]
    block_table_ptr, # int32 [num_requests, max_num_blocks_per_req]
    token_indices_ptr, # int32 [num_tokens, NUM_TOPK_TOKENS]
    cu_seqlens_ptr, # int32 [num_tokens + 1]
    out_ptr, # int32 [num_tokens, NUM_TOPK_TOKENS]
```

```

max_num_blocks_per_req: tl.constexpr,
BLOCK_SIZE: tl.constexpr,
BLOCK_N: tl.constexpr, # tile width along columns
bt_stride0, bt_stride1,
ti_stride0, ti_stride1,
):
    # 每个程序处理一个 token 的一列 tile
    token_id = tl.program_id(0)
    tile_id = tl.program_id(1)
    indice_id = tile_id * BLOCK_N + tl.arange(0, BLOCK_N)

    req = tl.load(req_id_ptr + token_id)
    seq_start = tl.load(cu_seqlens_ptr + token_id)
    seq_end = tl.load(cu_seqlens_ptr + token_id + 1)

    if tile_id * BLOCK_N + seq_start >= seq_end:
        return

    # 加载 token index 并计算 block_id 和 in-block offset
    ti_ptr = token_indices_ptr + token_id * ti_stride0 + indice_id * ti_stride1
    tok = tl.load(ti_ptr)
    is_invalid_tok = tok < 0
    block_id = tok // BLOCK_SIZE
    inblock_off = tok % BLOCK_SIZE

    valid_block = (block_id < max_num_blocks_per_req) & (block_id >= 0)
    bt_ptr = block_table_ptr + req * bt_stride0 + block_id * bt_stride1
    base = tl.load(bt_ptr, mask=valid_block, other=0)

    # 输出全局索引: block_id * BLOCK_SIZE + offset, 对 -1 标记置 0
    out_val = tl.where(is_invalid_tok | (~valid_block), 0, base * BLOCK_SIZE + inblock_off)
    out_ptr_ij = out_ptr + seq_start + indice_id
    out_ptr_ij_mask = (seq_start + indice_id) < seq_end
    tl.store(out_ptr_ij, out_val, mask=out_ptr_ij_mask)

```

vllm/model_executor/models/deepseek_v2.py

修改 Indexer forward 方法，为 ROCm 创建专用 RoPE 路径，避免影响其他平台。

```

# vllm/model_executor/models/deepseek_v2.py
# SparseAttnIndexer.forward 中的 ROCm 分支
if current_platform.is_rocm():
    # ROCm 专用路径: 合并 Q/K 拆分后直接调用 rotary_emb
    kw, _ = self.wk_weights_proj(hidden_states)
    k = kw[:, :self.head_dim]
    weights = kw[:, self.head_dim:]

    k = self.k_norm(k)

    # 此处直接使用 rotary_emb 的 in-place 版本，避免多余的 split/cat

```

```

    rotary_emb(positions, q[..., :self.rope_dim],
                k[..., :self.rope_dim].unsqueeze(1))
else:
    # 原始 CUDA 路径: 先 split 再 apply RoPE 再 cat
    q_pe, q_nope = torch.split(q, [self.rope_dim, self.head_dim - self.rope_dim], dim=-1)
    kw, _ = self.wk_weights_proj(hidden_states)
    k = kw[:, :self.head_dim]
    weights = kw[:, self.head_dim:]
    k = self.k_norm(k)
    k_pe, k_nope = torch.split(k, [self.rope_dim, self.head_dim - self.rope_dim], dim=-1)
    q_pe, k_pe = rotary_emb(positions, q_pe, k_pe.unsqueeze(1))
    q = torch.cat([q_pe.reshape(-1, self.n_head, self.rope_dim), q_nope], dim=-1)
    k = torch.cat([k_pe.squeeze(-2), k_nope], dim=-1)

```

评论区精华

1. 断言位置优化: gemini-code-assist 指出将 head_num 不支持检查放在 get_mla_padded_q 的热路径中可能影响性能, 建议移到配置初始化阶段。作者未回复但很可能接受, 因为在最终代码中该断言保留, 但后续 PR 可能调整。
 2. Indexer 页大小匹配: tjtanana 询问 indexer 的 page size 是否需要与 kernel block size 一致 (当前支持 [1,64]), ganyi1996ppo 确认 ROCm sparse MLA 不依赖 block size, 并建议保持与 NVIDIA 一致但也可支持 block-size 1。
 3. ROCm 特定路径: tjtanana 建议为 ROCm 创建单独路径而非修改通用代码, ganyi1996ppo 回应已用 current_platform.is_rocm() 分支隔离, 避免影响其他平台。
 4. Pre-commit 修复: tjtanana 提醒修复 pre-commit 问题, ganyi1996ppo 随后提交修复 commit。
- 断言位置是否应放在配置阶段 (design): 作者未直接回复, 但最终代码保留断言。后续可能改善。
 - Indexer block size 与 kernel block size 的匹配关系 (question): ganyi1996ppo 解释 ROCm sparse MLA 不依赖 block size, 建议保持与 NVIDIA 一致的 block size (64) 但也可添加 block-size 1 支持。最终 indexer 返回 [1,64]。
 - ROCm 特定路径 vs 修改通用路径 (design): 采用条件分支, 未创建独立文件。
 - Pre-commit 修复 (other): 已修复。

风险与影响

- 风险:
 1. 核心路径变更风险: 修改了 MLA attention backend 的核心 kernel 和 metadata 构建逻辑, 可能引入精度或稳定性问题。作者仅给出 GSM8K 分数, 未提供小模型微基准对比, 缺乏回归验证。
 2. FP8 稀疏 support 的兼容性: 新增的 FP8 decode 路径依赖 deepgemm_fp8_paged_mqa_logits 和特定 kernel, 若未覆盖所有 head 配置 (如 head=32 已明确不支持) 可能导致非法内存访问。

3. ROCm 平台隔离：在 `deepseek_v2.py` 中通过 `current_platform.is_rocm()` 分支，其他平台行为不变但仍需确认分支逻辑不影响 CUDA 后端。
4. 缺少测试配套：未提供相应的单元测试或集成测试文件，风险较高。历史 PR 中类似优化通常附有测试（如 `test_dcp_a2a.py`）。

- 影响：

1. 用户影响：ROCm 用户使用 DeepSeek V3.2 系列模型（含 MLA）的推理性能显著提升，并新增 FP8 KV cache 支持，降低显存占用。NVIDIA 用户无影响。
2. 系统影响：变更集中在 attention backend 的 MLA 相关文件，不影响其他注意力后端或模型；但 `deepseek_v2.py` 中 Indexer forward 路径的修改可能影响 CUDA 卸载路径。
3. 团队影响：需要维护新增的 Triton kernel 和 ROCm 分支代码；后续重构应确保通用性与平台专用代码的平衡。 - 风险标记：核心路径变更，缺少测试覆盖，ROCm 专用代码

关联脉络

- PR #32649（原始 PR 被本 PR 替换）：本 PR body 明确说明将替换 #32649。