

PR #43385 完整报告

vllm-project/vllm

[ROCm] [DSv4] [Perf] Support DeepSeek v4 MTP

合并时间: 2026-05-24 18:43

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/43385>

执行摘要

- 一句话: 为 DeepSeek V4 在 ROCm 上添加 MTP 推测解码支持
- 推荐动作: 值得精读, 特别是 MTP 实现 (mtp.py) 和稀疏注意力 prefill 优化 (rocm.py 中的 `combine_topk_swa_indices_kernel`)。这些设计展示了如何在 vLLM 架构中为特定硬件定制模型和算子, 对理解 vLLM 的推测解码和注意力机制实现有较高参考价值。

功能与动机

DeepSeek V4 需要 MTP (多令牌预测) 支持以实现推测解码加速, 而之前 AMD 路径仅通过符号链接依赖 NVIDIA 实现, 无法利用 ROCm 硬件特性 (如 Aiter、Triton 内核优化)。通过创建独立 AMD 模型文件 (`vllm/models/deepseek_v4/amd/`), 可以集成 ROCm 特定的内核 (如稀疏 prefill 索引组合、MegaMoE 输入 staging), 并解决 MTP 在 ROCm 上的精度和性能问题。PR body 中详细提供了 GSM8K 精度和吞吐量基准测试结果, 验证了正确性和加速效果。

实现拆解

1. 创建 AMD 专属模型实现: 新增 `vllm/models/deepseek_v4/amd/model.py` (+1612 行), 实现 `DeepseekV4MLP`、`DeepseekV4MegaMoEExperts` 等核心模块, 覆盖前向传播、MoE 路由、专家参数映射。原符号链接被删除。
2. 实现 MTP draft 模型: 新增 `vllm/models/deepseek_v4/amd/mtp.py` (+520 行), 实现 `DeepSeekV4MultiTokenPredictorLayer` 和 `DeepSeekV4MultiTokenPredictor`, 包含 V4 特有的 `e_proj/h_proj` 分离、HC 头压缩、权重加载重映射等。
3. 优化稀疏注意力 prefill: 修改 `vllm/models/deepseek_v4/amd/rocm.py`, 新增 `_combine_topk_swa_indices_kernel` Triton 内核和 `combine_topk_swa_indices` 函数, 将 top-k 索引和滑动窗口注意力索引合并为对齐的连续索引供 ROCm 稀疏注意力使用。
4. 调整 MoE 位矩阵元数据导入: 在 `gpt_oss_triton_kernels_moe.py` 的 `_patch_make_bitmatrix_metadata` 中, 对 ROCm 平台使用直接 `triton_kernels` 导入而非 `vllm.third_party`, 避免与 `site-packages` 版本冲突。
5. 注册新 attention metadata 类型: 在 `llm_base_proposer.py` 中, 将 `DeepseekV4ROCMAtterMLASparseMetadata` 和 `DeepseekV4ROCMAtterSparseSWAMetadata` 添加到 ROCm 允许的 attention 类型列表。

6. 基础设施调整：对 `rocm_aiter_mla_sparse.py` 进行控制流和配置键调整，以支持 V4 稀疏注意力。

关键文件：

- `vllm/models/deepseek_v4/amd/model.py`（模块 模型层；类别 source；类型 core-logic；符号 `DeepseekV4MLP`, `init`, `forward`, `_deepseek_v4_stage_mega_moe_inputs_kernel`）：新增的 AMD 专属 DeepSeek V4 模型实现，包含核心 MLP、MegaMoE 层等，替代了之前的符号链接，是这次变更的主体。
- `vllm/models/deepseek_v4/amd/mtp.py`（模块 模型层；类别 source；类型 core-logic；符号 `DeepSeekV4MultiTokenPredictorLayer`, `init`, `forward`, `DeepSeekV4MultiTokenPredictor`）：新增的 AMD 专属 MTP draft 模型实现，包含 `MultiTokenPredictorLayer` 和 `MultiTokenPredictor`，是推测解码的关键组件。
- `vllm/models/deepseek_v4/amd/rocm.py`（模块 注意力；类别 source；类型 core-logic；符号 `_combine_topk_swa_indices_kernel`, `combine_topk_swa_indices`）：修改的 ROCm 稀疏 prefill 优化，新增 Triton 内核用于合并 top-k 和 SWA 索引，是稀疏注意力的核心优化。
- `vllm/model_executor/layers/fused_moe/experts/gpt_oss_triton_kernels_moe.py`（模块 MoE；类别 source；类型 dependency-wiring）：修改 MoE 位矩阵元数据 kernel 的导入路径，根据平台区分 ROCm 和 CUDA 的 `triton_kernels` 源，修复编译时崩溃。
- `vllm/v1/spec_decode/llm_base_proposer.py`（模块 推测解码；类别 source；类型 dependency-wiring）：注册新的 DeepSeek V4 特定的 attention metadata 类型到 ROCm 允许列表，使 MTP draft 模型能正确使用这些后端。
- `vllm/v1/attention/ops/rocm_aiter_mla_sparse.py`（模块 注意力；类别 infra；类型 infrastructure）：对 ROCm 稀疏注意力操作进行控制流和配置键调整，以支持 DeepSeek V4 的稀疏注意力需求。

关键符号：`_combine_topk_swa_indices_kernel`, `combine_topk_swa_indices`, `DeepseekV4MLP.forward`, `DeepSeekV4MultiTokenPredictorLayer.forward`, `DeepSeekV4MTP.forward`, `load_weights`

关键源码片段

`vllm/models/deepseek_v4/amd/model.py`

新增的 AMD 专属 DeepSeek V4 模型实现，包含核心 MLP、MegaMoE 层等，替代了之前的符号链接，是这次变更的主体。

```
class DeepseekV4MLP(nn.Module):
    # DeepSeek V4 的 MLP 层，支持标准 TP 和序列并行模式。
    # 如果是序列并行，则输入输出在 tp 组内分片，权重复制，无需额外规约；
    # 否则使用标准 TP，最后做 allreduce。
    def __init__(
        self,
        hidden_size: int,
        intermediate_size: int,
        hidden_act: str,
```

```

        swiglu_limit: float | None = None,
        quant_config: QuantizationConfig | None = None,
        reduce_results: bool = True,
        is_sequence_parallel: bool = False,
        prefix: str = "",
    ) -> None:
        super().__init__()

        self.gate_up_proj = MergedColumnParallelLinear(
            hidden_size,
            [intermediate_size] * 2,
            bias=False,
            quant_config=quant_config,
        )
        self.down_proj = RowParallelLinear(
            intermediate_size,
            hidden_size,
            bias=False,
            quant_config=quant_config,
            reduce_results=reduce_results,
            is_sequence_parallel=is_sequence_parallel,
        )

        # 使用 SiLU 与 Mul 激活，可选 clamp（用于 fp4 精度控制）
        if swiglu_limit is None:
            self.act = SiluAndMul()
        else:
            self.act = SiluAndMulWithClamp(swiglu_limit)

    def forward(self, x):
        # 通过 gate_up_proj 生成 gate 和 up 的中间结果，激活后与 down_proj 相乘
        gate_up, _ = self.gate_up_proj(x)
        x = self.act(gate_up)
        x, _ = self.down_proj(x)
        return x

```

vllm/models/deepseek_v4/amd/mtp.py

新增的 AMD 专属 MTP draft 模型实现，包含 MultiTokenPredictorLayer 和 MultiTokenPredictor，是推测解码的关键组件。

```

class DeepSeekV4MultiTokenPredictorLayer(nn.Module):
    # DeepSeek V4 的 MTP 层，每个预测块包含一个 decoding layer 和 linear 投影。
    def __init__(
        self,
        vllm_config: VllmConfig,
        topk_indices_buffer: torch.Tensor,
        prefix: str,
        aux_stream_list: list[torch.cuda.Stream] | None = None,
    ) -> None:

```

```

super().__init__()

assert vllm_config.speculative_config is not None
config = vllm_config.speculative_config.draft_model_config.hf_config
self.config = config
quant_config = vllm_config.quant_config
self.rms_norm_eps = config.rms_norm_eps

# 创建 V4 decoding layer (内部包含 self-attention 和 MoE FFN)
self.decoder = DeepseekV4DecoderLayer(
    vllm_config=vllm_config,
    prefix=f"{prefix}.decoder",
    aux_stream_list=aux_stream_list,
)

# 线性投影: 从隐藏维到 2 * hidden_size (用于下一步输入的 emb 合成)
self.linear = ReplicatedLinear(
    config.hidden_size,
    2 * config.hidden_size,
    bias=False,
    quant_config=quant_config,
    prefix=f"{prefix}.linear",
)

```

vllm/models/deepseek_v4/amd/rocm.py

修改的 ROCm 稀疏 prefill 优化，新增 Triton 内核用于合并 top-k 和 SWA 索引，是稀疏注意力的核心优化。

```

# ROCm 稀疏 prefill 保持此密集合并本地化，因此 AMD 特定的 SWA 修改
# 不会触及共享的 DeepSeek V4 缓存工具。
_SPARSE_PREFILL_TOPK_ALIGNMENT = 128

```

```

@triton.jit
def _combine_topk_swa_indices_kernel(
    combined_indices_ptr, combined_indices_stride,
    combined_lens_ptr,
    topk_indices_ptr, topk_indices_stride,
    query_start_loc_ptr,
    seq_lens_ptr, gather_lens_ptr,
    M, N,
    TOP_K: tl.constexpr, COMPRESS_RATIO: tl.constexpr,
    WINDOW_SIZE: tl.constexpr, TOPK_WIDTH: tl.constexpr,
    PADDED_TOP_K: tl.constexpr,
):
    # 并行化: batch 维度作为 program_id(0)，内部 token 用 worker_id 拆分
    batch_idx = tl.program_id(0)
    worker_id = tl.program_id(1)
    num_workers = tl.num_programs(1)

```

```

base = tl.load(query_start_loc_ptr)
query_start = tl.load(query_start_loc_ptr + batch_idx) - base
query_end = tl.load(query_start_loc_ptr + batch_idx + 1) - base
query_len = query_end - query_start
seq_len = tl.load(seq_lens_ptr + batch_idx)
gather_len = tl.load(gather_lens_ptr + batch_idx)
start_pos = seq_len - query_len
gather_start = seq_len - gather_len

for token_idx in range(query_start + worker_id, query_end, num_workers):
    token_idx_in_query = token_idx - query_start
    pos = start_pos + token_idx_in_query
    topk_len = tl.minimum((pos + 1) // COMPRESS_RATIO, TOP_K)
    swa_len = tl.minimum(pos + 1, WINDOW_SIZE)

    # 写入 top-k 索引（已偏移 batch ID）
    topk_offset = tl.arange(0, PADDED_TOP_K)
    topk_mask = topk_offset < topk_len
    safe_topk_offset = tl.where(topk_offset < TOPK_WIDTH, topk_offset, 0)
    topk_indices = tl.load(
        topk_indices_ptr + token_idx * topk_indices_stride + safe_topk_offset,
        mask=topk_mask, other=-1,
    )
    valid_topk = (topk_indices >= 0) & (topk_indices < N)
    topk_indices = tl.where(valid_topk, topk_indices + M * batch_idx, -1)
    tl.store(combined_indices_ptr + token_idx * combined_indices_stride + topk_offset,
            topk_indices, mask=topk_mask)

    # 写入 SWA 索引（紧接在 top-k 之后）
    swa_offset = tl.arange(0, WINDOW_SIZE)
    tl.store(
        combined_indices_ptr + token_idx * combined_indices_stride + topk_len + swa_offset,
        M * batch_idx + N + swa_offset + pos - swa_len + 1 - gather_start,
        mask=swa_offset < swa_len,
    )
    tl.store(combined_lens_ptr + token_idx, topk_len + swa_len)

```

评论区精华

Review 中主要讨论点：

- 导入路径变更：zyongye 提问为何某些文件只修改了导入但没有实际引用。tjtanaa 解释整个 vllm 已改用直接导入 triton_kernels，旧路径会导致与 site-packages 版本不匹配并引发错误（如 tensor_details.bitmatrix 的编译时崩溃）。
- GDN bug 修复：tjtanaa 在评论中指出 rocm_aiter_fusion.py 中 GDN 导入路径错误，会导致 ModuleNotFoundError，并说明该修复已在 #43486 PR 中独立处理。
- 后续清理计划：tjtanaa 表示将在后续 PR 中删除 NVIDIA 文件夹中不再使用的 `forward_native` 路径（原为 ROCm 设置）。

- 导入路径变更原因 (question): 确认导入路径修改必要, tjtanaa 提供了错误日志示例。
- GDN 导入 bug 修复 (bugfix): 修复已独立存在, 本 PR 不包含。
- 后续清理 forward_native 路径 (design): 留作后续 PR 处理。

风险与影响

- 风险:
 1. 性能退化风险: 在高并发 (≥ 64) 下 MTP 模式吞吐量下降 37%, 用户若在高负载场景启用 MTP 可能反效果。需根据实际负载决策。
 2. 内存与精度风险: 使用 fp8_e4m3 缓存和 fp4 索引器, 可能引入数值精度问题。尽管 GSM8K 精度测试通过, 但其他 benchmark 需验证。
 3. 维护兼容性: 新增 AMD 专用模型文件后, 若 NVIDIA/XPU 需类似改动, 需保持接口一致。未来合并时可能冲突。
 4. 测试覆盖不足: 无自动化单元测试 (仅依赖手动 benchmark 和 lm-eval), 回归风险较高。
- 影响:
 - 用户: ROCm 用户使用 DeepSeek V4 时可通过 `--speculative_config` 启用 MTP, 在低并发场景获得显著吞吐提升。默认不启用, 不影响现有工作流。
 - 系统: 新增约 2.2K 行代码, 集中在 `deepseek_v4/amd/` 下, 模块化清晰, 对核心框架侵入小。
 - 团队: 确立了 AMD 专属模型文件的组织方式, 为后续平台特定优化提供了参考。
 - 风险标记: 高并发性能退化, 缺少自动化测试, FP8/FP4 精度风险, 平台兼容维护成本

关联脉络

- PR #43486 [ROCm][Critical] Fix the GDN import bug: 本 PR 中提及 GDN 导入 bug, 并在评论中指出已由 #43486 独立修复。
- PR #43142 [kv_offload]: Add DSv4 support: 同属 DeepSeek V4 功能支持, 可能与 MTP 的 KV 管理相关。