

PR #50585 完整报告

vllm-project/vllm

[K3 Perf] Optimize k3 dspark fused kv, 4.5~4.6x kernel performance improvement

合并时间: 2026-08-08 03:16

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/50585>

执行摘要

- 一句话: K3 DSpark 融合 KV 投影重构, 核函数提速约 4.5 倍
- 推荐动作: 值得精读。推荐关注三个设计点: 一是如何用 MergedColumnParallelLinear + shard_id 在权重加载期零拷贝实现跨层融合投影, 复用 vLLM 标准量化分发; 二是 precompute_and_store_context_kv 从「运行期拼权重 + 回退分支」收敛为「标准线性层 + 元数据」的过程, 对比 base/head 可看到手工拼权重代码的大幅删除; 三是 current_workspace_manager() 替代常驻 buffer 的显存管理方式。若团队计划支持 Kimi-K3-NVFP4 或后续量化 DSpark drafter, 此 PR 是直接前置。

功能与动机

PR body 明确说明目标: 旧实现是 `5 x full projections -> compute Q + KV -> discard Q`, 新实现是 `1 x fused projection -> compute only KV (fused kv)`, 核心动机是去掉冗余的 Q 投影计算、减少投影次数。review 中 tlrnchlsmth 进一步推动用标准 vLLM linear 层 + QuantKey 分发替代 K3 专属实现, 以便支持 HuggingFace 上 RedHatAI/Kimi-K3-NVFP4 及 benchislett 预告的后续量化 NVFP4 DSpark drafter。

实现拆解

1. 权重加载契约改造 (`vllm/models/kimi_k3/nvidia/dspark_mla.py`): 新增模块级函数 `_duplicate_context_kv_weights(weights, num_layers)`, 在 `K3DSparkForCausalLM.load_weights` 中先于 `AutoWeightsLoader.load_weights` 处理权重流。对每个名字含 `.self_attn.kv_a_proj_with_mqa.` 的权重, 在原样放行给每层 `self_attn` 的同时, 用 `weight.detach()` 复制一份并设置 `fused_weight.shard_id = layer_idx`, 映射到 `context_kv_proj.{param_name}`。`detach()` 保证零拷贝 (底层 `data_ptr` 共享), `shard_id` 让 `MergedColumnParallelLinear` 的加载器把各层权重按列拼接。这替代了 base 版本加载后从 `fused_qkv_a_proj.weight` 中 narrow 出 KV 列再 `torch.cat` 的手工拼权重逻辑, 使权重加载完全走标准线性层契约, 量化权重 (如 `weight_packed`、`weight_scale`) 也能自动分发。
2. 模块结构改造 (`K3DSparkModel.__init__`): 新增 `context_kv_proj`, 类型为 `MergedColumnParallelLinear(hidden_size, [kv_width] * num_hidden_layers, bias=False, return_bias=False, quant_config=self.quant_config, disable_tp=True)`, prefix 复用第一层 `layers.{start_layer_id}.self_attn.fused_qkv_a_proj`, `disable_tp=True` 保证每个 TP rank 持有完整副本、无需通信。`kv_width = kv_lora_rank +`

qk_rope_head_dim (K3 为 576) 。

3. 推理路径简化: precompute_and_store_context_kv 删除 _context_kv_fusion_available 状态机与量化回退分支 (base 版本在 quant_config 非空时逐层调用 fused_qkv_a_proj 并丢弃 Q 列), 现在仅在首次调用时通过 _build_fused_context_kv_metadata 收集各层 kv_a_layernorm.weight 与几何参数, 随后直接走 _precompute_fused_context_kv: 一次 context_kv_proj(context_states) 产出全部层 KV ((num_ctx, L, kv_width)), 拆出 kv_c 与 k_pe 后做跨层 grouped RMSNorm 与 RoPE 旋转。A 投影 FLOPs 从 5*2112 行降为 5*576 行, 减少约 72.7%。
4. 显存管理优化: 删除常驻的 _context_positions_repeated 预分配 buffer (原先按 max_num_batched_tokens 常驻显存), 改为在 _precompute_fused_context_kv 内通过 current_workspace_manager().get_simultaneous(...) 按需申请并复用 workspace 显存, 每次仅取前 num_layers * num_ctx 个元素使用。
5. 测试配套 (tests/models/test_dspark_mla.py) :
test_k3_dspark_uses_replicated_markov_head 增加 MergedColumnParallelLinear 的 recording dummy (make_context_kv_proj), 并为 config 补充 kv_lora_rank=3, qk_rope_head_dim=1、为 vllm_config 补充 scheduler_config.max_num_batched_tokens=16, 断言 context_kv_proj 以 (8, [4])、disable_tp=True、prefix 等参数正确构造; 新增 test_context_kv_weights_are_loaded_as_merged_linear_shards, 构造带 weight_packed、weight_scale 的权重流, 验证复制出的 context_kv_proj 权重与原权重 data_ptr 相同且 shard_id 正确。

关键文件:

- vllm/models/kimi_k3/nvidia/dspark_mla.py (模块 草稿模型; 类别 source; 类型 data-contract; 符号 _duplicate_context_kv_weights, _build_fused_context_kv_metadata, _build_fused_context_kv_buffers, _precompute_fused_context_kv) : 源码主路径: 新增 _duplicate_context_kv_weights 实现权重加载期零拷贝复制与 shard_id 打标, 新增 context_kv_proj (MergedColumnParallelLinear), 删除手工拼权重 buffer 与量化回退分支, 融合路径统一走 _precompute_fused_context_kv, 并用 current_workspace_manager 替代常驻 positions buffer。
- tests/models/test_dspark_mla.py (模块 模型测试; 类别 test; 类型 test-coverage; 符号 make_context_kv_proj, test_context_kv_weights_are_loaded_as_merged_linear_shards, test_k3_dspark_uses_replicated_markov_head) : 测试配套: 扩展 test_k3_dspark_uses_replicated_markov_head 断言 context_kv_proj 的构造参数 (输出维度、prefix、disable_tp), 并新增 test_context_kv_weights_are_loaded_as_merged_linear_shards 验证权重复制契约、shard_id 与 data_ptr 共享。

关键符号: _duplicate_context_kv_weights, _build_fused_context_kv_metadata, _precompute_fused_context_kv, precompute_and_store_context_kv, load_weights, make_context_kv_proj, test_context_kv_weights_are_loaded_as_merged_linear_shards, test_k3_dspark_uses_replicated_markov_head

关键源码片段

vllm/models/kimi_k3/nvidia/dspark_mla.py

源码主路径：新增 `_duplicate_context_kv_weights` 实现权重加载期零拷贝复制与 `shard_id` 打标，新增 `context_kv_proj` (`MergedColumnParallelLinear`)，删除手工拼权重 `buffer` 与量化回退分支，融合路径统一走 `_precompute_fused_context_kv`，并用 `current_workspace_manager` 替代常驻 `positions` `buffer`。

```
def _duplicate_context_kv_weights(
    weights: Iterable[tuple[str, torch.Tensor]], num_layers: int
) -> Iterable[tuple[str, torch.Tensor]]:
    """把每层 KV 投影权重同时喂给跨层融合的 context_kv_proj。"""
    for name, weight in weights:
        # 原始权重照常流向每层 self_attn 的 fused_qkv_a_proj
        yield name, weight
        # 只处理 MLA 的 KV 投影权重，识别标记形如
        # layers.3.self_attn.kv_a_proj_with_mqa.weight_packed
        layer_prefix, marker, param_name = name.partition(
            ".self_attn.kv_a_proj_with_mqa."
        )
        if not marker:
            continue
        layer_idx_str = layer_prefix.rsplit(".", 1)[-1]
        if not layer_idx_str.isdecimal():
            continue
        layer_idx = int(layer_idx_str)
        if layer_idx >= num_layers:
            continue
        # detach 后原地打 shard_id: MergedColumnParallelLinear 的加载器
        # 会依 shard_id 把各层权重拼接为对应输出列，且不复制底层存储
        fused_weight = weight.detach()
        fused_weight.shard_id = layer_idx
        yield f"context_kv_proj.{param_name}", fused_weight

def _precompute_fused_context_kv(
    self,
    context_states: torch.Tensor,
    context_positions: torch.Tensor,
    context_slot_mapping: torch.Tensor | list[torch.Tensor | None] | None,
) -> None:
    num_ctx = context_states.shape[0]
    num_layers = self._num_context_layers

    # 一次 KV-only GEMM 取代 5 次全量 Q+KV GEMM：对 K3 而言投影行数从
    # 5*2112 降为 5*576，A 投影 FLOPs 减少约 72.7%
    all_kv = self.context_kv_proj(context_states)
    all_kv = all_kv.view(num_ctx, num_layers, self._context_kv_width)
    all_kv_c = all_kv[..., : self._context_kv_lora_rank]
    all_k_pe = all_kv[..., self._context_kv_lora_rank :]

    # Layer-major 排布让 2-D RMSNorm 权重在单个 grouped kernel 里
```

```

# 为每个草稿层选中独立一行
all_kv_c = all_kv_c.permute(1, 0, 2).contiguous()
all_kv_c_normed = torch.empty_like(all_kv_c)
ops.rms_norm(
    all_kv_c_normed,
    all_kv_c,
    self._context_kv_norm_weights,
    self._context_rms_norm_eps,
)

all_k_pe = all_k_pe.permute(1, 0, 2).contiguous()
all_k_pe_flat = all_k_pe.view(num_layers * num_ctx, 1, self._context_rope_dim)
# 从 workspace 按需申请 positions buffer, 替代模块内常驻的
# _context_positions_repeated 大块显存
(repeated_positions,) = current_workspace_manager().get_simultaneous(
    ((num_layers * self._max_num_context_tokens,), torch.int64),
)
repeated_positions = repeated_positions[: num_layers * num_ctx]
repeated_positions.view(num_layers, num_ctx).copy_(context_positions)
# 后续 RoPE 旋转 (含 k_pe 与 1-D repeated_positions 对齐) 与
# do_kv_cache_update 写入逻辑与既有实现保持一致, 此处不展开

```

tests/models/test_dspark_mla.py

测试配套: 扩展 `test_k3_dspark_uses_replicated_markov_head` 断言 `context_kv_proj` 的构造参数 (输出维度、prefix、disable_tp), 并新增 `test_context_kv_weights_are_loaded_as_merged_linear_shards` 验证权重复制契约、`shard_id` 与 `data_ptr` 共享。

```

def test_context_kv_weights_are_loaded_as_merged_linear_shards():
    weights = [
        (
            "layers.0.self_attn.kv_a_proj_with_mqa.weight_packed",
            torch.arange(4),
        ),
        (
            "layers.1.self_attn.kv_a_proj_with_mqa.weight_scale",
            torch.tensor(0.5),
        ),
    ]

    duplicated = dspark_mla._duplicate_context_kv_weights(weights, 2)
    mapped = list(K3DSparkForCausalLM.hf_to_vllm_mapper.apply(duplicated))

    # 每个原始权重后跟一个复制到 context_kv_proj 的副本, 顺序正确
    assert [name for name, _ in mapped] == [
        "model.layers.0.self_attn.fused_qkv_a_proj.weight_packed",
        "model.context_kv_proj.weight_packed",
        "model.layers.1.self_attn.fused_qkv_a_proj.weight_scale",
        "model.context_kv_proj.weight_scale",
    ]

```

```
# 0 层副本 shard_id 为 0, 其余 (原始权重与 1 层副本) 为 1
assert [weight.shard_id for _, weight in mapped] == [1, 0, 1, 1]
# detach() 复制保证零拷贝: 副本与原权重共享底层存储
assert mapped[0][1].data_ptr() == mapped[1][1].data_ptr()
assert mapped[2][1].data_ptr() == mapped[3][1].data_ptr()
```

评论区精华

评审核心围绕「K3 专属融合逻辑是否应抽象为通用机制」展开。tlrmchlsmth 在 PR review 中建议: *Instead of just adding support for FP8 here, can we represent the GEMM as a normal vLLM linear layer and dispatch based on a QuantKey?*。作者最初倾向保持 K3 专属 (回应: *This 5 layer fusion is for k3 only, no other models will reuse the mechanism, so I am thinking keep it as it is now, and do the refactor when there is a similar structure?*) , 但考虑到要支持 HuggingFace 的 RedHatAI/Kimi-K3-NVFP4 变体以及 benchislett 预告的 *Quantized NVFP4 DSpark drafters are also coming eventually*, 最终采纳建议重构为 *MergedColumnParallelLinear* 版本, 并报告新版本性能进一步提升 (full speedup 从约 3.9 倍升至 4.5~4.6 倍)。另两条行内评论: 解引用 *weight* 前应检查字段是否存在 (针对量化层 *qweight* 场景) ——作者已修复; *positions* buffer 应由 *current_workspace_manager()* 分配——作者已修复。

- 是否用标准 linear 层 + QuantKey 替代 K3 专属手工拼权重 (design): 最终重构为 *MergedColumnParallelLinear* + *shard_id* 方案, 性能从约 3.9 倍提升至 4.5~4.6 倍并获得 APPROVED。
- 解引用 *proj.weight* 前需要检查字段存在性 (correctness): 作者回复 *Nice catch, solved* 并修复。
- *positions* buffer 改用 *current_workspace_manager* (performance): 已改为按需从 *workspace* 申请, 减少模型常驻显存。
- 量化 NVFP4 DSpark drafter 的后续支持 (question): 本 PR 的 *MergedColumnParallelLinear* 改造为量化 drafter 通过 QuantKey 分发提供了基础, 具体支持留待后续 PR。

风险与影响

- 风险:
 1. 量化路径覆盖变化: base 版本在 *quant_config* 非空时回退到逐层投影, 删除该回退后融合路径必须依赖 *MergedColumnParallelLinear* 的量化分发 (QuantKey 机制); 单元测试只覆盖 *quant_config=None* 的场景, 末端到端验证 FP8/NVFP4 量化 checkpoint 下 *kv_a_proj_with_mqa* 权重 (*weight_packed/weight_scale*) 经 *shard_id* 拼接后的数值正确性, 存在回归风险。
 2. 权重名解析契约: *_duplicate_context_kv_weights* 依赖 *self_attn.kv_a_proj_with_mqa*. 标记与层号前缀解析, 若 checkpoint 命名变化会静默跳过复制, *context_kv_proj* 权重缺失时缺少显式校验。
 3. v1 依赖引入: 模型文件新增 *from vllm.v1.worker.workspace import current_workspace_manager*, 模型层与 v1 worker 耦合, 非 v1 执行环境可能受影响。

4. 性能验证局限：4.5~4.6 倍数据来自作者的单机脚本（B300、单一 shape 组合），未做端到端吞吐验证，且作者明确表示缺少 GPU 资源；不同 GPU 与 num_ctx 下加速比可能浮动。 - 影响：影响范围集中在 K3 DSpark 推测解码路径：

precompute_and_store_context_kv 是上下文 KV 预计算的热点，本 PR 使其从 5 次全量投影降为 1 次 KV-only 投影，A 投影 FLOPs 减少 72.7%，B300 实测核函数提速 4.5~4.6 倍，可直接降低首 token 延迟；权重加载路径新增一次零拷贝迭代复制，开销可忽略；对 K3 以外的模型无行为影响（逻辑全部在 vllm/models/kimi_k3/nvidia/dspark_mla.py 内）。对团队而言，该 PR 把跨层融合投影收敛到标准

MergedColumnParallelLinear 契约上，降低了未来接入量化 DSpark drafter（如 NVFP4）的适配成本，但同时也让模型层依赖 v1 workspace，后续维护需注意该耦合。

- 风险标记：量化回退路径移除，端到端性能未验证，依赖 v1 workspace，权重名解析敏感

关联脉络

- PR #51253 [ROCm][Perf] Kimi-K3 Shard Latent MoE up-projection for ROCm path: 同属 Kimi-K3 性能优化系列（K3 模型路径），虽然平台不同（ROCm vs NVIDIA），但与本 PR 共同构成 K3 推理性能的持续优化脉络。