

PR #42864 完整报告

vllm-project/vllm

[ROCm][Compile] Fuse AR + RMSNorm + per-group FP8 quant (+ DSv3.2 indexer fan-out)

合并时间: 2026-06-09 20:06

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/42864>

执行摘要

- 一句话: 融合 AR+RMSNorm+per-group FP8 量化为单一 AITER 内核
- 推荐动作: 值得精读, 尤其是对 ROCm 编译优化感兴趣的读者。
- 关注点: 如何通过 VllmPatternReplacement 处理多 consumer 扇形输出, 使用 emit_bf16 复用 fused op 输出。
- 讨论中关于 MatcherQuantFP8 的演进和使用方式值得借鉴。
- 该 PR 是持续优化的一个环节, 展示了 vLLM 编译 fusion pass 的扩展性。

功能与动机

DSv3.2 等 FP8 块级量化模型的每个 transformer 块以 all_reduce -> [fused_add]rms_norm -> per-group FP8 quant -> fp8_gemm 结尾。PR #41825 修复了 rms_norm -> group_fp8_quant 的融合, 但当 rms_norm 已被 AR+RMS 吸收后, 独立的量化 kernel 仍然存在, 导致每 decode step 约 535μs 的开销。本 PR 通过将量化吸收进 AR epilogue 来消除该开销。

实现拆解

实现步骤如下:

1. 在 vllm/_aiter_ops.py 中注册两个新的 Custom op:
rocm_aiter_fused_allreduce_rmsnorm_quant_per_group (返回 FP8 量化输出、residual、scale) 和带 emit_bf16=True 的变体, 返回额外 bf16 归一化激活。新增 fused_ar_rms_per_group_quant 协议方法和 feature probe 函数。
2. 在 vllm/compilation/passes/fusion/allreduce_rms_fusion.py 中添加三个 VllmPatternReplacement 类:
 - AiterAllreduceFusedRMSNormGroupQuantFP8Pattern: 匹配 AR -> rms_norm -> group_fp8_quant (无 residual)
 - AiterAllreduceFusedAddRMSNormGroupQuantFP8Pattern: 匹配 AR -> fused_add_rms_norm -> group_fp8_quant (单 consumer)
 - AiterAllreduceFusedAddRMSNormGroupQuantWithIndexerPattern: 匹配 AR -> fused_add_rms_norm -> (group_fp8_quant + rocm_unquantized_gemm) (两个 consumer, DSv3.2 indexer 扇形输出) 所有模式使用 MatcherQuantFP8 进行量化匹配, 注册顺序确保最大子图优先。

3. 添加 `weight.to(input.dtype)` 转换以修复 AITER 内核的 dtype 约束（从先前修复 cherry-pick）。
4. 在 `tests/compile/passes/distributed/test_fusion_all_reduce.py` 中添加 `TestAiterAllReduceRMSNORMGroupQuantFP8Model` 和 `test_rocm_aiter_all_reduce_rmsnorm_group_quant_fp8_fusion_pass_replace` 测试，覆盖三个新模式的替换正确性。
5. 经过多轮 review 后，最终通过 rebase 消除了 `use_triton_quant` 标志，统一使用 `MatcherQuantFP8` 进行匹配。

关键文件：

- `vllm/_aiter_ops.py`（模块 Aiter 操作注册；类别 source；类型 core-logic；符号 `fused_ar_rms_per_group_quant`, `_rocm_aiter_fused_allreduce_rmsnorm_quant_per_group_impl`, `_rocm_aiter_fused_allreduce_rmsnorm_quant_per_group_fake`, `_rocm_aiter_fused_allreduce_rmsnorm_quant_per_group_with_bf16_norm_impl`）：注册 Custom op 和协议方法，是融合操作的底层实现基础。新增 `fused_ar_rms_per_group_quant` 协议和两个 custom op（含 bf16 变体），以及 feature probe。
- `vllm/compilation/passes/fusion/allreduce_rms_fusion.py`（模块 编译融合；类别 source；类型 core-logic；符号 `AiterAllreduceFusedRMSNORMGroupQuantFP8Pattern`, `AiterAllreduceFusedAddRMSNORMGroupQuantFP8Pattern`, `AiterAllreduceFusedAddRMSNORMGroupQuantWithIndexerPattern`, `init`）：核心模式匹配文件，新增三个 `VllmPatternReplacement` 类，是融合逻辑的主入口。
- `tests/compile/passes/distributed/test_fusion_all_reduce.py`（模块 融合测试；类别 test；类型 test-coverage；符号 `TestAiterAllReduceRMSNORMGroupQuantFP8Model`, `init`, `_group_quant`, `_dequantize_to_bf16`）：新增单元测试，验证三个新模式的替换正确性，确保融合后输出与原路径一致。

关键符号：`AiterAllreduceFusedRMSNORMGroupQuantFP8Pattern.init`,
`AiterAllreduceFusedRMSNORMGroupQuantFP8Pattern.pattern`,
`AiterAllreduceFusedRMSNORMGroupQuantFP8Pattern.replacement`,
`AiterAllreduceFusedAddRMSNORMGroupQuantFP8Pattern.init`,
`AiterAllreduceFusedAddRMSNORMGroupQuantFP8Pattern.pattern`,
`AiterAllreduceFusedAddRMSNORMGroupQuantFP8Pattern.replacement`,
`AiterAllreduceFusedAddRMSNORMGroupQuantWithIndexerPattern.init`,
`AiterAllreduceFusedAddRMSNORMGroupQuantWithIndexerPattern.pattern`,
`AiterAllreduceFusedAddRMSNORMGroupQuantWithIndexerPattern.replacement`,
`_rocm_aiter_fused_allreduce_rmsnorm_quant_per_group_impl`,
`_rocm_aiter_fused_allreduce_rmsnorm_quant_per_group_fake`,
`_rocm_aiter_fused_allreduce_rmsnorm_quant_per_group_with_bf16_norm_impl`,
`_rocm_aiter_fused_allreduce_rmsnorm_quant_per_group_with_bf16_norm_fake`,
`has_fused_allreduce_rmsnorm_quant_per_group`,
`TestAiterAllReduceRMSNORMGroupQuantFP8Model.init`,
`TestAiterAllReduceRMSNORMGroupQuantFP8Model.forward`,

TestAiterAllReduceRMSNormGroupQuantFP8Model._group_quant,
TestAiterAllReduceRMSNormGroupQuantFP8Model._dequantize_to_bf16,
test_rocm_aiter_all_reduce_rmsnorm_group_quant_fp8_fusion_pass_replace

评论区精华

Review 中核心讨论包括：

- MatcherQuantFP8 兼容性 (来自 [gemini-code-assist\[bot\]](#) 和 [ProExpertProg](#)) : 初始版本为 indexer 扇形输出模式使用直接 Triton op 引用, reviewer 建议改用 MatcherQuantFP8 以确保与现有模式一致。经作者尝试后发现原因为 MatcherQuantFP8 当时只跟踪 AITER 路径。ProExpertProg 提出 MatcherQuantFP8 已在 main 上修复支持两种路径, 作者通过 rebase 后成功统一, 去除了 use_triton_quant 标志。
- weight dtype 转换开销 (Rohan138) : 询问 weight.to(input.dtype) 是否会引入额外 kernel。作者确认该转换是从另一个 PR cherry-pick 的修复, 因为 AITER 内核要求激活和权重 dtype 一致。
- Feature probe 清理 TODO (Rohan138) : 请求添加 TODO 注释, 在 AITER 版本升级后移除 `has_fused_allreduce_rmsnorm_quant_per_group` 探测函数。作者提交 commit 添加了 TODO。
- MatcherQuantFP8 兼容性 & use_triton_quant 标志消除 (design): 通过 rebase 到最新 main, 使用单个 MatcherQuantFP8 实例即可匹配两路径, 消除了双注册需求。
- weight.to(input.dtype) 转换是否会引入额外 kernel (correctness): 该转换是必要的 dtype 适配, 预计不会显著影响性能 (通常权重已为目标 dtype 或仅需一次转换)。后续可通过优化避免。
- Feature probe 清理 TODO (testing): 已添加 TODO, 计划在 AITER 0.1.14 发布后清理。

风险与影响

- 风险：风险分析：
 1. 依赖 AITER 版本：新模式依赖于 ROCm/aiter#2823 引入的 `fused_ar_rms_per_group_quant` launcher。代码通过 `has_fused_allreduce_rmsnorm_quant_per_group()` 探测, 不存在时自动回退到旧的 AR+RMS-only 融合, 因此不会崩溃, 但性能回退到以前水平。
 2. 模式匹配范围：当前模式仅匹配 `group_shape=(1, 128)` 的 FP8 量化, 若其他模型使用不同的 group size 则不会触发。
 3. 测试覆盖：单元测试覆盖了三种模式, 但仅使用 torch.mm 而非真实 FP8 GEMM, 可能遗漏部署中的数值问题。
 4. 性能回退：若 AITER 内核出现问题 (如精度), 静默回退到旧融合, 性能下降但功能正常。
- 影响：影响分析：
 - 影响范围：仅 ROCm 平台 (MI355X), 且需要 AITER 库支持。受益模型主要为 DeepSeek V3.2 及其他使用 per-group FP8 量化和 indexer 的模型。

- 性能影响: DSV3.2 TP4 上 TPOT 从 17.69ms 降至 16.96ms (约 4%), 消除了每 decode step 的独立量化 kernel (约 535 μ s)。
- 代码影响: 新增约 795 行代码, 主要分布在 fusion pass 和 op 注册。不改变现有 Python API 或用户界面。
- 团队影响: ROCm 编译团队需要维护新模式, 但设计遵循已有模式架构, 理解成本低。
- 风险标记: 依赖 AITER 版本, 模式匹配范围有限, 仅 ROCm 平台

关联脉络

- PR #41825 [ROCm][Compile] Fuse RMSNorm + per-group FP8 quant (RocmAiterRMSNormQuantFusionPass): 本 PR 是 #41825 的 AR-side 对应部分, #41825 实现了 rms_norm -> group_fp8_quant 的融合, 本 PR 在此基础上进一步融合 all_reduce, 消除残留的独立量化 kernel。