

PR #46730 完整报告

vllm-project/vllm

[ROCm][Perf][Bugfix] DSv4 indexer: use platform FP8 dtype (fnuz) for Q-quant on gfx942

合并时间: 2026-07-01 17:33

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/46730>

执行摘要

- 一句话: 修复 DSv4 indexer Q 量化 FP8 dtype 不匹配
- 推荐动作: 值得精读, 尤其关注 Triton 内核 constexpr 参数化设计, 以及如何通过平台检测实现硬件适配。PR 改动简洁 (1 文件 +19/-8), 但性能收益巨大, 是 ROCm 生态的重要修复。

功能与动机

在 gfx942 上, DeepSeek-V4 Flash indexer 的 Q 量化硬编码为 e4m3fn, 而 K 缓存使用平台原生 e4m3fnuz, 导致 FP8 logits 内核降级到混合 dtype 路径, 每次调用均走 fallback。PR body 指出此变更使两者统一为 fnuz, 让 logits 内核运行原生路径, 显著提升预填充性能。

实现拆解

1. 导入平台检测: 在 fused_indexer_q.py 中增加 from vllm.platforms import current_platform, 用于获取当前平台的 FP8 dtype。
2. 修改内核常数: 向 Triton 内核 _fused_indexer_q_rope_quant_kernel 添加两个 constexpr 参数 FP8_MAX 和 USE_FNUZ, 分别控制量化最大值和 FP8 dtype 选择。当 USE_FNUZ=True 时, FP8_MAX=224.0, 存储 dtype 为 tl.float8e4b8; 否则 FP8_MAX=448.0, 存储 dtype 为 tl.float8e4nv。fnuz 最大值的选取与 quant_utils.py 中的 get_fp8_min_max() 一致。
3. 更新宿主函数逻辑: 在 fused_indexer_q_rope_quant 中, 通过 current_platform.fp8_dtype() 获取平台 dtype, 判断是否为 float8_e4m3fnuz, 据此设置 fp8_max 和 fp8_dtype 变量, 并传递给内核。同时将 index_q_fp8 的张量 dtype 改为平台 dtype, 而非硬编码的 float8_e4m3fn。
4. NVIDIA 和 MXFP4 路径不受影响: 两个新内核 constexpr 有默认值 (FP8_MAX=448.0, USE_FNUZ=False), 因此其他后端无需修改。

关键文件:

- vllm/models/deepseek_v4/common/ops/fused_indexer_q.py (模块 算子; 类别 source; 类型 bugfix; 符号 _fused_indexer_q_rope_quant_kernel, fused_indexer_q_rope_quant)
: 核心变更文件: 修复 Q 量化 FP8 dtype, 使 gfx942 上 Q 和 K 统一使用 fnuz, 启用原生 FP8 logits 路径。

关键符号: _fused_indexer_q_rope_quant_kernel, fused_indexer_q_rope_quant

关键源码片段

vllm/models/deepseek_v4/common/ops/fused_indexer_q.py

核心变更文件：修复 Q 量化 FP8 dtype，使 gfx942 上 Q 和 K 统一使用 fnuz，启用原生 FP8 logits 路径。

```
# vllm/models/deepseek_v4/common/ops/fused_indexer_q.py ( 关键片段 )
import torch
from vllm.platforms import current_platform
from vllm.triton_utils import tl, triton
# ...

def _fused_indexer_q_rope_quant_kernel(
    # ... 其他参数
    FP8_MAX: tl.constexpr = 448.0, # fnuz 用 224.0, ocp 用 448.0
    USE_FNUZ: tl.constexpr = False, # gfx942 上 True
):
    # ... 计算 amax
    index_q_scale = tl.div_rn(tl.maximum(amax, 1e-4), FP8_MAX)
    index_q_scale = tl.math.exp2(tl.math.ceil(tl.math.log2(index_q_scale)))
    # 选择 fp8 dtype: fnuz 用 tl.float8e4b8, ocp 用 tl.float8e4nv
    fp8_dtype = tl.float8e4b8 if USE_FNUZ else tl.float8e4nv
    # 使用 fp8_dtype 存储
    tl.store(fp8_base_ptr + nope_offset, tl.div_rn(x_nope, index_q_scale).to(fp8_dtype))
    # ...

def fused_indexer_q_rope_quant(...):
    # ...
    fp8_dtype = current_platform.fp8_dtype() # 获取平台 fp8 dtype
    use_fnuz = fp8_dtype == torch.float8_e4m3fnuz
    fp8_max = 224.0 if use_fnuz else 448.0
    index_q_fp8 = torch.empty_like(index_q, dtype=fp8_dtype) # 使用平台 dtype
    # ... 调用内核时传递 FP8_MAX=fp8_max, USE_FNUZ=use_fnuz
```

评论区精华

审核者 tjtanad 要求提供端到端模型 lmeval 分数以验证正确性。作者 akii96 补充了 GSM8K 20-shot 结果：基线 exact_match 为 0.9227 (flexible-extract)，PR 后为 0.9224，差异在误差范围内，确认无精度退化。

- 验证正确性：要求提供 lmeval 分数 (testing): akii96 提供了 GSM8K 20-shot 结果，基线 0.9227 与 PR 后 0.9224 差异在误差范围内，验证通过。

风险与影响

- 风险：低风险。变更仅影响 gfx942 平台 (is_fp8_fnuz() == True)，其他平台 (gfx950、NVIDIA) 因默认 constexpr 值不受影响。但需注意：当前只通过一个 Triton 内核测试 (test_fused_indexer_q_rope_quant.py)，9/10 形状完美匹配，一个形状因 RoPE 舍入差

异有少量误差，非 dtype 问题。建议在更多负载下验证，确保 fnuz 量化路径与下游 logits 内核兼容。

- 影响：正面影响显著：gfx942 上预填充 TTFT 最高加速 7.4 倍，精度无损。影响范围限于 ROCm gfx94x 平台上的 DeepSeek-V4 Flash 模型，其他平台无变化。团队可期待 ROCm 推理性能大幅提升。
- 风险标记：平台特定修复，测试覆盖不全

关联脉络

- PR #41601 [ROCM] DeepSeek-V4 Flash attention enablement (stalled): 同为 ROCm DSv4 enablement PR，包含同样的 indexer dtype 修复，但范围更大，因 rebase 停滞。此 PR 提取了最小修复。
- PR #42033 [ROCM] DeepSeek-V4 Flash attention enablement (stalled): 同上，是另一 stalled PR。
- PR #43950 [ROCM][DSV4] Use aiter mHC pre/post as the default ROCm path: 同为 ROCm DSv4 性能优化 PR，涉及同一模型系列。