

PR #50458 完整报告

vllm-project/vllm

[Kimi K3 Bug] Fix deepgemm support for kimi k3

合并时间: 2026-07-31 08:10

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/50458>

执行摘要

- 一句话: 修复 Kimi K3 下 DeepGEMM BLOCK_D 断言崩溃
- 推荐动作: 改动虽小但值得精读: 它以最小变更修复了一个真实的高 TP 模型崩溃, 揭示了 DeepGEMM EP 路径对 hidden_size 的隐含对齐约束; math.gcd 求块大小的做法也可作为类似 kernel 对齐问题的通用解法。若后续继续扩展 DeepGEMM 支持的非 1024 约数形状, 建议顺手排查 deep_gemm_utils.py 中其他同类 min(...) 计算。

功能与动机

PR body 给出了完整复现命令: `VLLM_USE_RUST_FRONTEND=0 vllm serve moonshotai/Kimi-K3 --moe-backend auto --tensor-parallel-size 8 ...`, 启动时在 `deep_gemm_utils.py` 第 428 行 `assert hidden_size % BLOCK_D == 0` 处抛出 `AssertionError`。根因是 `ep_gather` 的 Triton kernel 块大小 `BLOCK_D` 固定为 `min(hidden_size, 1024)`, 隐含假设 `hidden_size` 是 1024 的约数; Kimi K3 的 latent-MoE 在 TP=8 下 `hidden_size=384`, 不满足该约束。

实现拆解

1. 变更入口: `vllm/model_executor/layers/fused_moe/deep_gemm_utils.py`, 文件顶部新增 `import math`。
2. 核心逻辑修复: 在 `ep_gather` 函数中把 `BLOCK_D = min(hidden_size, 1024)` 改为 `BLOCK_D = math.gcd(hidden_size, 1024)`。这样任意 `hidden_size` 都能选出不超过 1024 且可整除的块大小, 原有的 `assert` 保留为兜底; 对 DeepSeek V4 等 4096/2048 形状, gcd 结果与原 min 相同, 行为完全不变。
3. 测试配套: `tests/kernels/moe/test_deepgemm.py` 在 `FP4_MNKs` 参数化列表中新增 (128, 384, 3584), 即 Kimi K3 TP8 latent-MoE 的 FP4 形状, 使其进入 `test_deepgemm_fp4_vs_triton` 的覆盖范围, 并对 DeepGEMM 输出与 BF16 参考做 <5% 误差校验。无配置、schema 或部署配套改动。

关键文件:

- `vllm/model_executor/layers/fused_moe/deep_gemm_utils.py` (模块 MoE 内核; 类别 source; 类型 core-logic; 符号 `ep_gather`): 修复的核心文件: `ep_gather` 的 `BLOCK_D` 从 `min(hidden_size, 1024)` 改为 `math.gcd(hidden_size, 1024)`, 是 Kimi K3 崩溃点的直接修复。

- tests/kernels/moe/test_deepgemm.py (模块 内核测试; 类别 test; 类型 test-coverage; 符号 FP4_MNKs) : 在 FP4_MNKs 参数化中补充 Kimi K3 TP8 latent-MoE 形状, 覆盖 hidden_size=384 这一非 1024 约数的回归场景。

关键符号: ep_gather

关键源码片段

vllm/model_executor/layers/fused_moe/deep_gemm_utils.py

修复的核心文件: ep_gather 的 BLOCK_D 从 min(hidden_size, 1024) 改为 math.gcd(hidden_size, 1024), 是 Kimi K3 崩溃点的直接修复。

```
import math
```

```
import torch
```

```
# 其余导入省略, 核心修复集中在 ep_gather 的块大小计算
```

```
@torch.no_grad()
```

```
def ep_gather(
```

```
    input_tensor: torch.Tensor,
```

```
    recv_topk_ids: torch.Tensor,
```

```
    recv_topk_weight: torch.Tensor,
```

```
    input_index: torch.Tensor,
```

```
    expert_map: torch.Tensor | None,
```

```
    output_tensor: torch.Tensor,
```

```
):
```

```
    num_warps = 2
```

```
    num_tokens = output_tensor.shape[0]
```

```
    hidden_size = input_tensor.shape[1]
```

```
    # 修复前: BLOCK_D = min(hidden_size, 1024), 要求 hidden_size 能被 1024 整除,
```

```
    # 否则 assert 直接崩溃 (Kimi K3 在 TP=8 下 latent-MoE 的 hidden_size=384) 。
```

```
    # 修复后: 取最大公约数, 任意 hidden_size 都能找到不超过 1024 的可整除块大小。
```

```
    BLOCK_D = math.gcd(hidden_size, 1024)
```

```
    assert hidden_size % BLOCK_D == 0
```

```
    grid = (triton.cdiv(hidden_size, BLOCK_D), min(num_tokens, 1024))
```

```
# 启动 Triton kernel, BLOCK_D 作为编译期常量传给 _fwd_kernel_ep_gather
```

```
_fwd_kernel_ep_gather(grid)(
```

```
    num_tokens,
```

```
    input_tensor,
```

```
    input_tensor.stride(0),
```

```
    input_tensor.stride(1),
```

```
    recv_topk_ids,
```

```
    recv_topk_ids.stride(0),
```

```
    recv_topk_ids.stride(1),
```

```
    recv_topk_weight,
```

```
    recv_topk_weight.stride(0),
```

```
    recv_topk_weight.stride(1),
```

```

    input_index,
    input_index.stride(0),
    input_index.stride(1),
    output_tensor,
    output_tensor.stride(0),
    output_tensor.stride(1),
    topk_num=recv_topk_ids.shape[1],
    expert_map=expert_map,
    HAS_EXPERT_MAP=expert_map is not None,
    num_warps=num_warps,
    BLOCK_D=BLOCK_D,
)
return

```

tests/kernels/moe/test_deepgemm.py

在 FP4_MNKs 参数化中补充 Kimi K3 TP8 latent-MoE 形状，覆盖 hidden_size=384 这一非 1024 约数的回归场景。

```

# FP4 测试参数化覆盖：新增 Kimi K3 TP8 latent-MoE 形状 (128, 384, 3584),
# 其中 hidden_size=384 不是 1024 的约数，专门覆盖本次 ep_gather 的块大小修复。
FP4_MNKs = [
    (128, 4096, 4096), # DeepSeek V4 shape
    (256, 2048, 2048), # Half-size variant
    (128, 384, 3584), # Kimi-K3 TP8 latent-MoE shape
]

```

评论区精华

该 PR 未产生实质性技术讨论。[sfeng33](#) 直接 APPROVED 且未附评论；[claude\[bot\]](#) 仅推送了仓库自动 review 订阅提示 ("Comment [@claude review](#) for a one-time review")。结论：修复方案直观，维护者认可，无遗留未解决疑虑。

- 暂无高价值评论线程

风险与影响

- 风险：

1. 回归风险：math.gcd(hidden_size, 1024) 对原有受支持形状（hidden_size 是 1024 的约数，如 4096、2048）结果与原 min 完全一致，无行为变化；仅对非约数形状改用更小的 BLOCK_D，可能带来轻微 kernel 效率变化，但正确性受 assert 保护。
2. 覆盖风险：修复只针对 ep_gather 一处；deep_gemm_utils.py 中若存在其他基于 min(hidden_size, 1024) 的假设，本次未一并排查，可能在其他非约数形状下再次暴露。
3. 环境风险：新增测试依赖 DeepGEMM 支持环境 (is_deep_gemm_supported())，无该环境的 CI 会自动 skip，回归守护能力有限。- 影响：影响范围：仅 DeepGEMM 专家并行 (EP) 路径的 ep_gather，即 Kimi K3 在 NVIDIA 平台启用 DeepGEMM 的 serving 场景。对用户：修复了 Kimi K3 在 TP=8 长上下文 (max-model-len 1048576) 场景下的启动崩溃，使该模型可正常服务。对系统：无配置、接口或部署变化。对团队：把

latent-MoE 的非标准 hidden_size 纳入测试矩阵，降低后续内核重构再次破坏该形状的风险。 - 风险标记：MoE 内核路径变更，依赖 DeepGEMM 支持环境，仅修复单点块大小计算

关联脉络

- PR #50516 [ROCm][CI] Fall back to lossless Kimi K3 MXFP4 emulation on gfx942: 同为 Kimi K3 的 MXFP4/DeepGEMM 支持修复，改动 mxfp4.py 与 fused_moe 专家内核路径，与本 PR 处于同一功能线。
- PR #50242 K3 DSpark AR fusion: 同为 Kimi K3 模型在 fused_moe 内核层的能力增强，且与 deep_gemm 工具函数同属 vllm/models/kimi_k3 的推理路径。