

PR #51682 完整报告

vllm-project/vllm

[Bugfix][Kimi-K3] Give the AMD packed KDA decode kernel the state-index stride

合并时间: 2026-08-11 00:05

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/51682>

执行摘要

- 一句话: 修复 AMD KDA 解码内核缺少的 state 索引步长
- 推荐动作: 值得精读, 尤其是两点设计决策: 其一是坚持把 stride 支持做进 kernel 而不是在调用方做 contiguous 拷贝, 避免 decode 热路径的分配与 memcpy; 其二是用 PACKED_DECODE_IMPLS 映射把 NVIDIA/AMD 两份 vendored 实现统一参数化, 让同一套语义测试同时约束双副本。对维护多后端 vendored kernel 的团队是很好的参考模式。

功能与动机

PR body 明确描述了问题链路: AMD 副本的 `fused_recurrent_kda_packed_decode` 假设 `state_indices` 连续且 unit-stride, 并在入口抛 `ValueError: state_indices must be contiguous and one-dimensional`; 而 NVIDIA 副本已支持 `stride_state_indices` 且只要求一维。差异在 speculative decoding 开启时才暴露——`GDNAttentionMetadataBuilder` 以 `block_table_tensor[:, 0]` 作为 KDA state slot, Mamba group 的 block table 宽度为 `1 + num_speculative_blocks` 列, 因此该列是 strided view, warmup 即失败。作者选择修复 kernel 而非调用方, 理由是 strided view 合法, 调用侧强制 contiguous 会在 decode 热路径上每步引入一次 device 分配和 memcpy。

实现拆解

1. 内核参数与寻址修复 (`vllm/models/kimi_k3/amd/ops/third_party/kda/fused_recurrent.py`): `fused_recurrent_kda_packed_decode_kernel` 新增 `stride_state_indices` 参数, 索引计算从 `state_indices + i_n` 改为 `state_indices + i_n * stride_state_indices`; 入口校验从 `ndim != 1 or stride(0) != 1` 放宽为仅 `ndim != 1`; kernel 启动处传入 `state_indices.stride(0)`。这样与 NVIDIA 副本完全对齐, decode 热路径不引入额外分配或拷贝。
2. 测试覆盖补齐 (`tests/models/kimi_k3/test_kda.py`): 新增 AMD 副本导入并别名, 构建 `PACKED_DECODE_IMPLS = {"nvidia": ..., "amd": ...}`; `test_packed_kda_decode_correctness` 增加 `impl` 参数化, 调用点改为 `PACKED_DECODE_IMPLS[impl]`。测试原本已参数化 `state_indices_stride ∈ [1, 8]`、`num_seqs ∈ [1, 8, 32]`、`lower_bound ∈ [-5.0, None]`, 但只覆盖 NVIDIA 副本, 这正是 AMD 缺失 stride 支持长期未被发现的原因; 修改后两份实现共用同一组参考与断言。
3. 验证与协作: 第二 commit 为合并 main 分支; CI 由维护者 `tjtanaa` 以 `/ci run` 触发并最终 approve。ROCm 7.2.4 / MI355X 上 24 个用例通过; 作者同时在 8x MI355X 上端到

端验证了 PP8/TP1、DP2xTP4 + EP、fp8 KV cache + speculative decoding 的 warmup 恢复，GSM8K 5-shot 保持 ~0.96。需注意端到端验证跑在带额外 Kimi-K3 工作的下游分支，单测结果才是本 PR 树上的证据。

关键文件：

- vllm/models/kimi_k3/amd/ops/third_party/kda/fused_recurrent.py（模块 KDA 内核；类别 source；类型 core-logic；符号 fused_recurrent_kda_packed_decode_kernel, fused_recurrent_kda_packed_decode）：本次修复的核心文件：AMD vendored KDA 解码内核补充 stride_state_indices 参数并按步长寻址，同时放宽入口校验，使 AMD 副本与 NVIDIA 副本行为收敛，直接消除 spec decode 场景的 warmup 失败。
- tests/models/kimi_k3/test_kda.py（模块 KDA 测试；类别 test；类型 test-coverage；符号 test_packed_kda_decode_correctness）：测试覆盖从仅 NVIDIA 副本扩展到 NVIDIA/AMD 双副本：通过 PACKED_DECODE_IMPLS 参数化，使 state_indices_stride 的既有参数矩阵同时约束两份 vendored 实现，防止同类漂移再次发生。

关键符号：fused_recurrent_kda_packed_decode,

fused_recurrent_kda_packed_decode_kernel, test_packed_kda_decode_correctness

关键源码片段

vllm/models/kimi_k3/amd/ops/third_party/kda/fused_recurrent.py

本次修复的核心文件：AMD vendored KDA 解码内核补充 `stride_state_indices` 参数并按步长寻址，同时放宽入口校验，使 AMD 副本与 NVIDIA 副本行为收敛，直接消除 spec decode 场景的 warmup 失败。

AMD vendored 的 KDA 解码内核。

修复前：state_indices 按 unit-stride 读取，并在入口拒绝任何非连续张量；

修复后：与 NVIDIA 副本对齐，支持任意步长的 1D 索引。

```
def fused_recurrent_kda_packed_decode_kernel(
    ...,
    stride_state_token: tl.constexpr,
    stride_state_indices, # 新增：state_indices 的步长，按运行时值传入
    H: tl.constexpr,
    K: tl.constexpr,
    V: tl.constexpr,
    ...,
):
    ...
    # 以 i_n * stride_state_indices 定位第 i_n 行的 state 索引，
    # 兼容 block_table_tensor[:, 0] 这类 strided view。
    state_idx = tl.load(state_indices + i_n * stride_state_indices).to(tl.int64)
    ...
```

```
def fused_recurrent_kda_packed_decode(...):
    ...
```

```

# 只要求 1D, 不再要求 stride(0) == 1:
# strided view 是合法输入; 若在调用方强制 contiguous 拷贝,
# 每次 decode 都会新增一次 device 分配与 memcopy。
if state_indices.ndim != 1:
    raise ValueError("`state_indices` must be one-dimensional.")
...
fused_recurrent_kda_packed_decode_kernel[grid](
    ...,
    stride_state_token=initial_state.stride(0),
    stride_state_indices=state_indices.stride(0), # 实际步长传给 kernel
    H=H,
    K=K,
    V=V,
)

```

tests/models/kimi_k3/test_kda.py

测试覆盖从仅 NVIDIA 副本扩展到 NVIDIA/AMD 双副本: 通过 `PACKED_DECODE_IMPLS` 参数化, 使 `state_indices_stride` 的既有参数矩阵同时约束两份 vendored 实现, 防止同类漂移再次发生。

```

# AMD 与 NVIDIA 的 KDA 内核是各自 vendored 的副本, 允许分叉;
# 因此共享语义的测试必须同时覆盖两份实现, 避免“一份支持 stride、
# 另一份不支持”的差异再次静默漂移。

```

```

PACKED_DECODE_IMPLS = {
    "nvidia": fused_recurrent_kda_packed_decode,
    "amd": fused_recurrent_kda_packed_decode_amd,
}

```

```

@pytest.mark.parametrize("num_seqs", [1, 8, 32])
@pytest.mark.parametrize("lower_bound", [-5.0, None])
@pytest.mark.parametrize("state_indices_stride", [1, 8])
@pytest.mark.parametrize("impl", PACKED_DECODE_IMPLS.keys())
@torch.inference_mode()
def test_packed_kda_decode_correctness(
    num_seqs: int,
    lower_bound: float | None,
    state_indices_stride: int,
    impl: str,
):
    ...
    # 两份实现共用同一套参考输入与误差断言,
    # 任何一份实现对于 state_indices 步长语义的回归都会在此暴露。
    packed_out, _ = PACKED_DECODE_IMPLS[impl](
        mixed_qkv=mixed_qkv,
        raw_g=raw_g,
        raw_beta=raw_beta,
        A_log=A_log,
    )

```

```
dt_bias=dt_bias,  
lower_bound=lower_bound,  
initial_state=packed_state,  
state_indices=state_indices,  
)
```

评论区精华

评审者 Fangzhou-Ai 评论“LGTM, but we need an unit test.”——认可修复方向，但明确要求补单测；第二个 commit 即把 `test_packed_kda_decode_correctness` 参数化到 AMD 副本，满足该要求。mergify[bot] 提示 pre-commit 初次失败，要求本地运行 `pre-commit run --all-files` 后提交；分支合并 main 后检查通过。维护者 tjtanaa 以 `/ci run` 触发 Buildkite CI (#83152) 并 approve。Claude 审核机器人因 fork PR 自动跳过审查，未产生额外讨论。

- 要求补充单元测试 (testing): 第二个 commit 将测试参数化到 `PACKED_DECODE_IMPLS` (nvidia/amd)，保留原有参数矩阵，AMD 副本的 stride 语义自此被单测约束。
- pre-commit 检查失败 (style): 分支合并 main 后检查通过，CI 触发成功。
- CI 触发与合入 (other): CI 通过后合入 main。

风险与影响

- 风险：行为变化仅限 AMD vendored 内核：索引由 unit-stride 改为按 stride(0) 寻址，对 1D 张量是正确的；但校验放宽后，若未来传入负步长或非常规 stride 的 1D view，kernel 不会显式拒绝，可能产生错误索引（当前唯一上游来源 `block_table_tensor[:, 0]` 是正步长列视图，风险低）。测试文件无条件导入 AMD vendored 模块，若该模块在非 ROCm 平台不可导入，可能引入新的 CI 失败——PR 未说明该模块的跨平台导入行为，属已知不确定性。decode 热路径无新增 allocation/memcpy，kernel 多一个运行时标量参数，开销可忽略；impl 参数化使测试矩阵翻倍，CI 耗时小幅上升。端到端验证不在本 PR 精确树上，复现性需注意。
- 影响：用户影响：AMD (MI355X 等 gfx950) 上 Kimi-K3 开启 speculative decoding 时的 warmup 崩溃被修复；未开启 spec decode 的场景行为不变。系统影响：decode 热路径无额外分配与拷贝，性能开销可忽略。团队影响：共享测试同时约束 NVIDIA/AMD 两份 vendored KDA 实现，防止双副本再次漂移，为 Kimi-K3 ROCm 后端长期维护设下护栏。
- 风险标记：vendored 双副本漂移，校验放宽，AMD 专有路径，端到端验证不在本 PR 树

关联脉络

- PR #51011 [ROCm][MLA] [K3] Fix fp8 KV cache decode on the AITER MLA backend: 同为 ROCm 上 Kimi-K3 解码路径 (AITER MLA / fp8 KV cache) 的 bugfix，与本 PR 同属 Kimi-K3 ROCm 后端稳定性加固线。
- PR #40958 [ROCm][CI] Extend ROCm AITER MHA (FA) coverage: 扩展 ROCm 注意力测试覆盖，与本 PR 用共享测试约束 vendored kernel 一致性的思路一脉相承。