

PR #45931 完整报告

vllm-project/vllm

[ROCm][DSV4] Disable TileLang MHC dispatch on gfx942

合并时间: 2026-06-22 17:26

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/45931>

执行摘要

- 一句话: 禁用 gfx942 上的 TileLang MHC 路径
- 推荐动作: 建议阅读: 本 PR 是典型的正确性保护修复, 代码逻辑清晰, 适合作为平台条件编译的参考模式。值得关注的设计决策: 采用装饰器函数 `_has_tilelang_mhc` 集中管理平台兼容性逻辑, 避免散落在各个 kernel 分支中; 区分 CUDA 和 ROCm 平台, 并为未来按 kernel 粒度控制预留可能。建议跟进: 后续可参考作者在 issue 中的 kernel 级分析, 按需重新启用 `mhc_post_tilelang` 和 `hc_head_fused_kernel_tilelang` 在 gfx942 上的使用, 以恢复部分性能。

功能与动机

关联 Issue #45698 报告 DeepSeek V4 TileLang MHC 路径在 gfx942 上产生错误输出。PR body 说明 "The non-TileLang fallback path produces coherent output, so gfx942 should not silently route through the known-incorrect TileLang MHC kernels until that path is fixed." 这是一个正确性保护措施。

实现拆解

1. 在 `vllm/model_executor/layers/mhc.py` 中添加平台检查函数 `_has_tilelang_mhc()`: 该函数首先判断是否安装了 TileLang, 然后区分 CUDA 和 ROCm 平台。对于 CUDA 始终允许 TileLang MHC; 对于 ROCm, 则调用 `on_gfx942()` 检测, 如果为 gfx942 则返回 `False`, 否则返回 `True` (例如 gfx950 仍启用)。全局常量 `HAS_TILELANG_MHC` 替换原有的 `HAS_TILELANG`。
2. 将 MHC 所有 kernel 分支 (`mhc_pre`、`mhc_post`、`hc_head`、`mhc_fused_post_pre`) 中的条件从 `HAS_TILELANG` 更换为 `HAS_TILELANG_MHC`: 确保在 gfx942 上所有 TileLang MHC 路径都被跳过, 回退到 `forward_native` 方法。
3. 更新 DeepSeek V4 模型文件: 在 `vllm/models/deepseek_v4/amd/model.py` 和 `vllm/models/deepseek_v4/amd/mtp.py` 中, 将导入从 `has_tilelang()` 改为 `HAS_TILELANG_MHC`, 并将 `self.has_tilelang` 属性改为赋值 `HAS_TILELANG_MHC`。
4. 调整 MHC kernel 测试: 在 `tests/kernels/test_mhc_kernels.py` 中, 将 `pytest.mark.skipif` 条件从 `current_platform.is_cuda_alike()` and `has_tilelang()` 替换为 `HAS_TILELANG_MHC`, 因此测试在 gfx942 上会自动跳过 TileLang MHC 相关用例, 避免误报失败。

关键文件：

- `vllm/model_executor/layers/mhc.py` (模块 模型执行器；类别 `source`；类型 `data-contract`；符号 `_has_tilelang_mhc`, `HAS_TILELANG_MHC`)：核心修改：引入 `_has_tilelang_mhc` 函数和 `HAS_TILELANG_MHC` 常量，替换原有 `HAS_TILELANG`，并在所有 `kernel` 分支 (`mhc_pre`, `mhc_post`, `hc_head`, `mhc_fused_post_pre`) 的 `forward_hip` 方法中将条件切换为新常量。
- `vllm/models/deepseek_v4/amd/model.py` (模块 模型定义；类别 `source`；类型 `data-contract`；符号 `has_tilelang`)：模型定义文件：导入 `HAS_TILELANG_MHC` 替代 `has_tilelang` 函数，并将 `self.has_tilelang` 属性赋值改为使用新常量。
- `vllm/models/deepseek_v4/amd/mtp.py` (模块 模型定义；类别 `source`；类型 `data-contract`；符号 `has_tilelang`)：MTP 模块文件：类似 `model.py` 的变更，导入 `HAS_TILELANG_MHC` 并赋值给 `self.has_tilelang`。
- `tests/kernels/test_mhc_kernels.py` (模块 MHC 测试；类别 `test`；类型 `test-coverage`)：测试文件：将所有 `skipif` 条件从 `current_platform.is_cuda_alike()` and `has_tilelang()` 更换为 `HAS_TILELANG_MHC`，从而在 `gfx942` 上自动跳过 `TileLang MHC` 测试用例。

关键符号：`_has_tilelang_mhc`, `MHCPreOp.forward_hip`, `MHCPostOp.forward_hip`, `HCHeadOp.forward_hip`, `MHCFusedPostPreOp.forward_hip`

评论区精华

Issue 评论中用户 @crazyguitar 反馈禁用 `TileLang MHC` 后吞吐量显著下降（从 ~1400 toks/s 降至更低），并附上一个针对 `torch` 后端的性能补丁。作者 @tuukkjs 随后进行了更精细的 `kernel` 分析，发现 `mhc_post_tilelang` 和 `hc_head_fused_kernel_tilelang` 在 `gfx942` 上正确，而 `mhc_pre_tilelang` 和 `mhc_fused_post_pre_tilelang` 全部失败。这暗示未来可以按 `kernel` 粒度恢复正确的路径，以缓解性能下降。但当前 PR 采取保守的全局禁用策略，优先保证正确性。

- 禁用 `TileLang MHC` 后的性能影响 (performance): 当前 PR 采取全局禁用策略以保证正确性，性能恢复留待后续按 `kernel` 优化。
- `TileLang MHC kernel` 在 `gfx942` 上的正确性细分 (correctness): 当前 PR 禁用所有 `TileLang MHC kernel` 是保守但正确的选择，后续可选择性启用通过的部分。

风险与影响

- 风险：
 1. 性能风险：在 `gfx942` 上回退到 `torch/triton` 实现后，`MHC kernel` 性能可能显著下降，影响 `DeepSeek V4` 模型吞吐量。Issue 评论中已有明确反馈。但这是当前唯一能保证正确性的选择。
 2. 遗漏其他 GPU 型号：`_has_tilelang_mhc` 默认对未列出的 `ROCm GPU` 返回 `False`，可能导致某些支持 `TileLang MHC` 的 `GPU` 也被禁用。不过当前已知仅 `gfx942` 有问题，且代码显式只对 `gfx942` 禁用，其他 `GPU on_gfx942()` 返回 `False` 则允许。
 3. 依赖 `vllm.platforms.rocm.on_gfx942`：该函数可能随 `ROCm` 库版本变化，但属于稳定 API。

4. 测试覆盖：测试跳过条件改为 HAS_TILELANG_MHC，在 gfx942 上 TileLang 测试被跳过，可能无法捕获 TileLang MHC 在 gfx942 上的未来修复。但这是有意为之。 - 影响：影响范围：仅限于 ROCm gfx942 (MI300X) 上的 DeepSeek V4 模型推理。CUDA 和 gfx950 不受影响。影响程度：严重影响：gfx942 上的 TileLang MHC 路径被完全禁用，可能导致性能下降（根据评论反馈显著），但确保了输出正确性。用户若不升级此修复，将面临错误输出。团队影响：AMD ROCm 团队需要后续修复 TileLang MHC kernel 以重新启用。

- 风险标记：性能退化，GPU 型号遗漏风险，测试跳过可能掩盖回归

关联脉络

- 暂无明显关联 PR