

PR #32230 完整报告

sgl-project/sglang

[AMD] MiniMax-M3: opt-in custom/quick all-reduce on ROCm

合并时间: 2026-07-30 17:56

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/32230>

执行摘要

- 一句话: MiniMax-M3 可选启用 custom/quick all-reduce
- 推荐动作: 建议精读。此 PR 展示了如何通过一个简单的条件判断和 opt-in 环境变量, 在不破坏现有兼容性的前提下, 为特定硬件平台释放性能优化。值得关注的是其设计权衡 (速度 vs 精度) 和最小侵入式实现方式。

功能与动机

MiniMax-M3 在 ROCm 上默认使用 NCCL all-reduce, 因为模型覆盖逻辑强制禁用了 custom all-reduce。而 sglang 已经内置了 **QuickAllReduce** (基于 mk1-project/quickreduce), 但 M3 无法使用。通过 opt-in 环境变量, 用户可以启用 INT4 quick-reduce, 在 MI355x TP4 上获得 8-16% 的吞吐提升和 17-19% 的 TTFT 降低。

实现拆解

1. 新增环境变量 **SGLANG_M3_ALLOW_CUSTOM_AR**: 在 python/sglang/srt/envron.py 的 Envs 类中添加 SGLANG_M3_ALLOW_CUSTOM_AR = EnvBool(False), 默认关闭。
2. 修改模型覆盖逻辑: 在 python/sglang/srt/arg_groups/overrides.py 的 _minimax_m3_overrides 函数中, 将原有的无条件禁用 custom all-reduce 改为条件判断: 仅当 aiter_fusion_resolved 为 False 且 SGLANG_M3_ALLOW_CUSTOM_AR 未设置时, 才设置 overrides["disable_custom_all_reduce"] = True。

关键文件:

- python/sglang/srt/arg_groups/overrides.py (模块 模型覆盖; 类别 source; 类型 core-logic; 符号 _minimax_m3_overrides): 核心逻辑: 修改 MiniMax-M3 的模型覆盖函数, 使 custom all-reduce 的禁用变为条件性, 受 SGLANG_M3_ALLOW_CUSTOM_AR 控制。
- python/sglang/srt/envron.py (模块 环境配置; 类别 source; 类型 core-logic): 定义新的 opt-in 环境变量 SGLANG_M3_ALLOW_CUSTOM_AR, 默认关闭, 确保向后兼容。

关键符号: _minimax_m3_overrides

关键源码片段

[python/sglang/srt/arg_groups/overrides.py](#)

核心逻辑：修改 MiniMax-M3 的模型覆盖函数，使 custom all-reduce 的禁用变为条件性，受 `SGLANG_M3_ALLOW_CUSTOM_AR` 控制。

```
# python/sglang/srt/arg_groups/overrides.py
# 在 _minimax_m3_overrides 函数中，ROCM 分支的处理逻辑

aiter_fusion_resolved = server_args.enable_aiter_allreduce_fusion
if (
    server_args.ep_size > 1
    and server_args.moe_a2a_backend == "none"
    and aiter_fusion_resolved
):
    logger.warning(
        "Disable --enable-aiter-allreduce-fusion for MiniMax-M3 "
        "standard EP on ROCm because the deferred fused all-reduce "
        "corrupts sparse MoE partial outputs."
    )
    overrides["enable_aiter_allreduce_fusion"] = False
    aiter_fusion_resolved = False

# 原本为：if not aiter_fusion_resolved:
# 现在额外检查 opt-in 环境变量，保持默认 NCCL 行为
if not aiter_fusion_resolved and not envs.SGLANG_M3_ALLOW_CUSTOM_AR.get():
    overrides["disable_custom_all_reduce"] = True
```

python/sglang/srt/envron.py

定义新的 opt-in 环境变量 `SGLANG_M3_ALLOW_CUSTOM_AR`，默认关闭，确保向后兼容。

```
# python/sglang/srt/envron.py
# 在 Envs 类的环境变量定义区域，新增以下条目

# MiniMax-M3 on ROCm force-disables custom all-reduce in its model override
# (arg_groups/overrides.py) when aiter all-reduce fusion is off. Set this to
# opt back in and keep custom/quick all-reduce enabled -- e.g. to run the
# INT4 quick-reduce path via ROCM_QUICK_REDUCE_QUANTIZATION={INT4,INT6,INT8}.
SGLANG_M3_ALLOW_CUSTOM_AR = EnvBool(False)
```

评论区精华

Review 中无实质性讨论，HaiShaw 直接批准了 PR。

- 暂无高价值评论线程

风险与影响

- 风险：风险较低。变更仅为 opt-in，默认行为不变。启用 custom all-reduce 后，INT4 量化会引入约 0.30 pp 的精度损失（GSM8K 验证），但用户可通过 `ROCM_QUICK_REDUCE_QUANTIZATION=INT6/INT8` 调整。无新 kernel 引入，回退路径明确。

- 影响：影响范围有限：仅影响使用 MiniMax-M3 模型且在 ROCm 平台上设置了 SGLANG_M3_ALLOW_CUSTOM_AR=1 的用户。默认情况下无任何变化。对 CUDA 和非 M3 模型无影响。
- 风险标记：opt-in 变更，精度损失风险（INT4）

关联脉络

- PR #32036 [ROCm] Add block-fp8 support for MiniMax-M3: 同为 MiniMax-M3 在 ROCm 上的优化，可叠加使用。
- PR #32190 [ROCm] FlyDSL-MoE for MiniMax-M3: 同为 MiniMax-M3 在 ROCm 上的优化，可叠加使用。