

PR #1835 完整报告

THUDM/slime

Add support for NVIDIA DGX Spark (GB10 / sm_121a, arm64)

合并时间: 2026-04-21 15:42

原文链接: <http://prhub.com.cn/THUDM/slime/pull/1835>

执行摘要

- 一句话: 新增 NVIDIA DGX Spark (GB10) 的 Docker 构建支持, 扩展硬件兼容性到 arm64 架构。
- 推荐动作: 该 PR 值得精读, 特别是对于需要在非 x86 架构或新兴硬件上部署的用户。关注 Dockerfile 中的环境变量设置和补丁机制, 展示了解决跨平台兼容性的实用技术权衡。

功能与动机

根据 PR body, GB10 是当前支持矩阵中的明确差距: slime 的发布镜像仅支持 x86_64, 且现有 arm64 基础镜像的 CUDA 12.9 缺少 sm_121a 目标, 导致 Triton JIT 崩溃。需要基于 CUDA 13.2 的 NGC vLLM 容器来提供最小可行基础。

实现拆解

1. 新增 Dockerfile.gb10 (docker/Dockerfile.gb10): 基于 NGC vLLM 容器 (arm64, CUDA 13.2), 设置环境变量如 TORCH_CUDA_ARCH_LIST, 安装系统依赖 (如 libz3-dev), 并应用头文件补丁 (如 NVTX 和 cuda_profiler_api.h)。
2. 添加内核编译补丁 (docker/patch/gb10/patch_sgl_kernel.py 和 sgl-kernel-arch.patch): 引入 SGL_KERNEL_GB10_ONLY CMake 选项, 仅编译 sm_120a 和 sm_121a 变体, 避免多架构编译时的内存溢出。
3. 修复代码语法错误 (slime/utils/arguments.py): 移除 --log-reward-category 参数 help 字符串中的多余逗号, 防止 argparse 解析时抛出 AttributeError。
4. 补充测试和文档: 新增 smoke 测试脚本 (scripts/run-qwen2.5-0.5B-gb10-smoke.sh) 用于快速验证, 并添加详细说明文档 (docker/NOTES_GB10.md) 记录 15 个阻塞问题的解决方案。

关键文件:

- docker/Dockerfile.gb10 (模块 部署脚本; 类别 infra; 类型 infrastructure): 新增的核心 Dockerfile, 定义了基于 NGC vLLM 容器的 arm64 构建流程, 解决了 GB10 的 CUDA 和 架构兼容性问题。
- docker/patch/gb10/patch_sgl_kernel.py (模块 补丁脚本; 类别 infra; 类型 infrastructure): Python 脚本用于动态修改 sgl-kernel 的 CMakeLists.txt, 添加 GB10 专用编译选项以减少内存使用。

- `slime/utils/arguments.py` (模块 参数解析; 类别 `source`; 类型 `core-logic`; 符号 `add_wandb_arguments`) : 修复了全局参数解析中的语法错误, 该错误会导致所有平台的 `--help` 调用失败, 影响用户体验。
- `scripts/run-qwen2.5-0.5B-gb10-smoke.sh` (模块 测试脚本; 类别 `other`; 类型 `core-logic`) : 新增的 `smoke` 测试脚本, 用于快速验证 GB10 环境下的 GRPO 训练流程是否正常工作。
- `docker/NOTES_GB10.md` (模块 文档说明; 类别 `docs`; 类型 `documentation`) : 详细文档记录了 15 个阻塞问题的根因和解决方案, 为后续维护和问题排查提供参考。

关键符号: `add_wandb_arguments`

关键源码片段

`docker/Dockerfile.gb10`

新增的核心 Dockerfile, 定义了基于 NGC vLLM 容器的 arm64 构建流程, 解决了 GB10 的 CUDA 和架构兼容性问题。

```
# 基于 NGC vLLM 容器 (arm64, CUDA 13.2) , 提供 GB10 所需的最小环境
FROM nvcr.io/nvidia/vllm:26.03-py3@sha256:
13e327dad79e6e417f6687fec2ba76b0386d597082ec0ee003c1e964ec6ad0e7

# 设置 CUDA 架构列表: GB10 使用 sm_121a, 但 NGC PyTorch 仅提供 compute_120
PTX, 通过前向兼容运行
ENV TORCH_CUDA_ARCH_LIST="12.0+PTX"

# 安装系统依赖: libz3-dev 用于 tilelang 的 Z3 SMT 调度器
RUN apt-get update && apt-get install -y --no-install-recommends \
    libz3-dev \
    && rm -rf /var/lib/apt/lists/*

# 应用头文件补丁: NVTX 和 cuda_profiler_api.h 在 CUDA 13 中被移除, 提供兼容性 shim
RUN git clone --depth 1 https://github.com/NVIDIA/NVTX.git /tmp/nvtx_src \
    && cp -r /tmp/nvtx_src/c/include/nvtx3 /usr/local/cuda/include/nvtx3 \
    && rm -rf /tmp/nvtx_src
COPY docker/patch/gb10/cuda_profiler_api.h /usr/local/cuda/include/cuda_profiler_api.h
```

`slime/utils/arguments.py`

修复了全局参数解析中的语法错误, 该错误会导致所有平台的 `--help` 调用失败, 影响用户体验。

```
def add_wandb_arguments(parser):
    # ... 其他参数定义 ...
    parser.add_argument(
        "--log-reward-category",
        type=str,
        default=None,
        help=(
            "Log statistics of the category of reward, such as why the reward function considers it as failed. "
            "Specify the key in the reward dict using this argument." # 修复: 移除末尾逗号, 确保
```

```
        help 为字符串而非元组
    ),
)
# ... 继续其他参数 ...
return parser
```

评论区精华

Issue 评论中，作者 boots-coder 分享了端到端 smoke 测试结果和 100-rollout 收敛实验的评估数据，验证了 GB10 支持的有效性；维护者 zhuzilin 表示肯定。无 review 评论，表明变更已通过测试验证并直接合并。

- 端到端测试验证 (testing): 测试验证了 GB10 支持的有效性，变更被认可并合并。

风险与影响

- 风险：技术风险包括：
 1. 兼容性风险：依赖特定 CUDA 13.2 版本和 NGC 容器，可能与其他环境或未来升级冲突；
 2. 维护复杂度：新增 Dockerfile 和补丁文件需要长期维护，可能增加构建流程的脆弱性；
 3. 性能不确定性：arm64 架构和 sm_121a 目标在 slime 中的性能表现尚未全面基准测试。
 - 影响：对用户影响：扩展了 slime 的硬件支持范围，使拥有 NVIDIA DGX Spark (GB10) 的用户能够运行训练流程。对系统影响：新增 arm64 构建路径，但不影响现有 x86_64 部署。对团队影响：需要维护新 Dockerfile 和补丁，但通过文档和测试降低了上手门槛。
 - 风险标记：新架构支持，依赖特定版本，构建复杂度增加

关联脉络

- PR #1813 [conda] Add install custom sgl-router to build_conda.sh: 同样涉及 Docker 和依赖配置的修改，展示了仓库中基础设施扩展的常见模式。
- PR #1849 fix: 修改脚本文件路径，与本 PR 的测试脚本添加相关，反映脚本管理的持续优化。