

PR #51208 完整报告

vllm-project/vllm

[ROCm][AMD][Installation] add LMCache kv-connector installation and runtime packages to docker image

合并时间: 2026-08-18 05:58

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/51208>

执行摘要

- 一句话: ROCm 镜像内置 LMCache, 发布默认启用并做架构裁剪
- 推荐动作: 值得精读。该 PR 虽不涉及核心推理代码, 但 Docker 多阶段构建的缓存正确性设计非常扎实: 用 ARG 拼 stage 名实现条件构建、digest-only COPY 作为缓存键、llvm-objdump --offloading 校验 code object 覆盖、sccache 契约复用。这些模式对任何“把第三方组件塞进镜像”的场景都有直接借鉴价值, 建议作为镜像工程化的模板案例。

功能与动机

PR body 明确说明: CUDA 的 `docker/Dockerfile` 在 `INSTALL_KV_CONNECTORS=true` 时会安装 LMCache, 且发布流水线对每个 `vllm/vllm-openai` tag 都开启; 但 `docker/Dockerfile.rocm` 从未引用该参数, 导致 `vllm/vllm-openai-rocm` 镜像不带 LMCache, 用户必须自行构建安装。本 PR 的目标就是让 ROCm 发布镜像与 CUDA 对齐, 开箱即用地提供 LMCache KV connector。

实现拆解

1. 构建开关与阶段选择: `docker/Dockerfile.rocm` 新增独立的 ARG `INSTALL_LMCACHE=false`, 并将最终阶段拆成 `final_lmcache_false` / `final_lmcache_true` 两个平行 stage, 通过 FROM `final_lmcache_${INSTALL_LMCACHE}` AS final 按开关选择。之所以用布尔值而非列表, 是因为 BuildKit 的 stage 名拼接要求值直接出现在 FROM 中; 若用列表则每次构建都会执行该 stage。
2. 源码获取与依赖修正: 新增 `lmcache_source` stage, 将 LMCache 锁定到 v0.5.3 对应的 commit (`140819c9...`), 回应了 bot 关于可编辑 tag 的供应链风险建议; 克隆后校验 HEAD 与引用一致, 并移除 `cufile-python`、`nvtx` 两个 NVIDIA-only 依赖, 把 `numpy<=2.2.6` 放宽到 `numpy<=2.3.5` 以匹配 release 镜像的实际版本。
3. wheel 构建与产物校验: `build_lmcache` stage 在 `BUILD_WITH_HIP=1 CXX=hipcc` 下用 `setup.py bdist_wheel` 编译 HIP 扩展, 支持 `USE_SCCACHE=1` 时启用 `/opt/sccache-wrappers`; 构建后用 `llvm-objdump --offloading` 校验 `c_ops*.so` 实际携带的 gfx 架构与目标集合完全一致, 防止“名义继承、实际缺架构”的静默降级。
4. 安装编排与层缓存: 发布路径优先下载 LMCache GitHub Releases 的预编译 wheel, `--no-index` 避免 pip 静默回退到 PyPI 上的 CUDA wheel; 当请求架构超出 wheel 覆盖 (

如 `gfx90a`) 时回退源码编译。wheel 本体不出现在最终层里——安装层只 `COPY` 一个约 118 B 的 `wheel.sha256` 作为缓存键, 实际 wheel 通过 bind mount 注入, 避免 7.8 MB blob 被 whiteout 残留在层里; 同时安装被放在 vLLM wheel 之前, 使稳定的 connector 层不被源码变动连带失效。

5. 发布与 CI 配套: `.buildkite/release-pipeline.yaml` 的 ROCm 发布构建显式传入 `--build-arg INSTALL_LMCACHE=true`; `.buildkite/scripts/ci-bake-rocm.sh` 把 `lmcache_source`、`build_lmcache` 纳入 CI base 的 stage 清单。验证环节包括镜像内 `lmcache server` 启动 (ZMQ/HTTP 均正常)、connector 注册列表、numpy 版本断言与 `uv pip check`。

关键文件:

- `docker/Dockerfile.rocm` (模块 镜像构建; 类别 infra; 类型 infrastructure; 符号 `lmcache_source`, `build_lmcache`, `final_lmcache_true`, `final_lmcache_false`): 核心变更文件, 新增 LMCACHE 的源码获取、HIP 编译、架构校验、安装编排全链路, 并决定镜像层缓存结构。
- `.buildkite/release-pipeline.yaml` (模块 发布流水线; 类别 config; 类型 configuration): ROCm 发布流水线显式开启 `INSTALL_LMCACHE=true`, 是“开箱即用”真正生效的入口。
- `.buildkite/scripts/ci-bake-rocm.sh` (模块 CI 脚本; 类别 other; 类型 core-logic; 符号 `DEFAULT_CI_BASE_DOCKERFILE_STAGES`): CI base 的 stage 清单加入 `lmcache_source` 和 `build_lmcache`, 直接影响 CI 缓存键与 base 镜像内容。

关键符号: `lmcache_source`, `build_lmcache`, `final_lmcache_true`, `final_lmcache_false`, `final_lmcache_${INSTALL_LMCACHE}`

评论区精华

review 中最有价值的交锋集中在四处:

- PYTHONHASHSEED 与 TP>1 哈希键一致性: jamesETsmith 建议在镜像里写死 `PYTHONHASHSEED=0`, 避免多 worker 哈希不一致导致缓存永远 miss; hongxiayang 反驳说任何固定 seed 都行, 且不同 seed 能隔离多租户共享 LMCACHE 后端的 keyspace, 不应把 0 烤进镜像, 最终只加注释提醒用户固定 seed。
- 缓存键正确性: AndreasKaratzas 指出 wheel 只通过 bind mount 进入安装层时, `LMCACHE_REF` 或 arch 变化后 warm builder 可能复用旧 install 层, 并引用 vLLM 线上事故 #44795 要求照搬 CUDA 的 `wheel.sha256` 模式; hongxiayang 进一步把 wheel 内容本身放进缓存键并加 `sha256sum -c` 二次校验。
- --no-deps 与依赖一致性: AndreasKaratzas 认为 `--no-deps` 会留下 knowingly inconsistent 的环境, 要求用上游认可的依赖集并跑 `uv pip check`; hongxiayang 论证去掉 `--no-deps` 会让 `numpy<=2.2.6` 降级 vLLM 自身依赖、破坏性更隐蔽, 并给出端到端实测 (`lmcache server` 启动、connector 注册、numpy/otel 版本不变)。
- 架构覆盖: AndreasKaratzas 指出非空默认 `LMCACHE_ROCM_ARCH="gfx942;gfx950"` 会让“继承镜像 `PYTORCH_ROCM_ARCH`”的 fallback 永远不生效, 且 `gfx90a` (MI200) 没有官方 wheel; 最终方案改为发布 wheel 优先 + 源码编译 fallback, 缺失架构打日志跳过。

- PYTHONHASHSEED 与 TP>1 哈希键一致性 (performance): hongxiayang 认为任何固定 seed 都有效, 且不同 seed 的 keyspace 可隔离多租户共享后端, 不应把 0 烤进镜像; 改为在 Dockerfile 中加注提示用户固定 seed。
- INSTALL_LMCACHE 独立 arg 避免缓存键污染 (design): hongxiayang 改为独立 ARG INSTALL_LMCACHE=false, 发布 job 显式传 true, 彻底解耦缓存键, 同时移除 umbrella 参数。
- wheel digest 作为安装层缓存键 (correctness): hongxiayang 实现 digest-only 导出 stage, 安装层只 COPY 约 118 B 的 digest, wheel 通过 bind mount 注入, 并加 sha256sum -c 二次校验。
- --no-deps 与依赖版本冲突 (numpy/otel) (correctness): 最终保留 wheel 直装 + numpy 版本断言 + uv pip check 兜底, 并在 PR body 中列出完整版本对比, 建议上游修正 requirements 上限。
- 架构覆盖: gfx90a 缺失与默认 arch 集 (design): 改为发布 wheel 优先 + 源码编译 fallback; 有效架构集与上游 wheel 取交集, 缺失 arch 打日志跳过, 无支持的 arch 集时构建失败。gfx90a 被明确排除并文档化。
- wheel 7.8 MB 进入镜像层的白化问题 (performance): 新增 digest-only 导出 stage, 安装层从 7.82 MB 降到 118 B, wheel 不在最终层中; hongxiayang 用隔离 harness 验证了 wheel 内容变化会强制重跑安装层。

风险与影响

• 风险:

1. 依赖版本冲突被绕过而非解决: --no-deps 使 LMCache 声明的 numpy<=2.2.6、opentelemetry-api<=1.40.0、opentelemetry-exporter-prometheus<=0.61b0 与镜像实际版本 (2.3.5、1.44.0、0.65b0) 不一致, 最终靠 uv pip check 和实测兜底, 但上游 requirements 与实测环境仍存在漂移风险。
1. gfx90a (MI200) 用户不可用: 发布镜像默认只携带 gfx942;gfx950 的 LMCache 设备代码, MI200用户拿到镜像后LMCache无法运行, 仅有构建日志提示, 缺少运行时文档说明。
2. 镜像体积增加: 整体 +921 MB (+2.19%), 其中 CuPy 约 366 MB, 对磁盘和拉取时间有实际影响。
3. 构建逻辑复杂度高: arch 字符串三种分隔符归一化、FROM final_lmcache_\${INSTALL_LMCACHE} 条件 stage、digest 缓存键、code-object 断言, 后续维护者需要同时理解 BuildKit 层缓存和 LMCache 上游发布形态, 否则容易引入静默错误。
4. 上游 wheel 覆盖变化: 若 LMCache 后续发布覆盖更多 arch, 发布流水线的下载路径会自动改变行为, release 镜像的 arch 集合会随之漂移, 需要配套监控。- 影响: 对用户: ROCm 镜像 (MI300X/MI325X、MI350X/MI355X 用户) 开箱即用地获得 LMCache KV connector, 省去手动编译安装; TP>1 用户需注意 PYTHONHASHSEED 固定。

对系统: 镜像体积 +2.19%; 构建层缓存结构显著优化, 避免 wheel blob 残留在层里、避免 vLLM 源码变动连带失效 connector 层。

对团队：新增与 LMCache 上游 release 的耦合点（arch 清单、版本 pin、requirements 修改），需要定期复核；同时为后续其他 KV connector 入镜像提供了可复用的构建模式。

- 风险标记：依赖版本冲突 (numpy/otel), gfx90a 无 LMCache, 镜像体积 +921 MB, --no-deps 绕过依赖解析，构建缓存键复杂度，上游 wheel 覆盖 arch 变化

关联脉络

- PR #44284 Relax CuPy constraint to only exclude 14.1.0: 同为 kv_connector 在 ROCm 侧的依赖治理，本 PR 安装 cupy-rocm-7-0 时也涉及 CuPy 版本约束处理。
- PR #52216 Promote prefix_cache_retention_interval to an argument and change the default to 0: 同属 kv-connector 功能线，说明 vLLM 正在把 KV 缓存外部化 / 多级化做实，本 PR 是 ROCm 镜像侧的基础设施配套。