

PR #48438 完整报告

vllm-project/vllm

[Bugfix] Preserve Marlin runtime tensor storage across weight reload

合并时间: 2026-07-31 05:50

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/48438>

执行摘要

- 一句话: 修复 Marlin 内核工作区和排序索引在权重重载时地址失效
- 推荐动作: 建议量化相关开发者精读, 了解 Marlin 内核的工作区模式和权重重载契约。FP8/NVFP4/MXFP4 等回退路径的同步修复模式值得学习。其他开发者可跳过。

功能与动机

Marlin 内核在权重重载时重新分配工作区和排序索引张量, 但层拷贝恢复只保护注册的参数 / 缓冲区, 导致 CUDA 图捕获后引用过期设备指针, 引发错误或错误输出。该问题是 RFC #48312 'Weight Reload Correctness for RL' 中定义的高风险项之一 (见 PR body)。

实现拆解

1. 修改 `marlin_make_workspace_new` (`vllm/model_executor/layers/quantization/utils/marlin_utils.py`): 新增 `existing` 可选参数, 当传入兼容的已有工作区时零填充并返回原存储; 不兼容时抛出 `ValueError`。
2. 修改 `MarlinLinearKernel.process_weights_after_loading` (`vllm/model_executor/kernels/linear/mixed_precision/marlin.py`): 调用时传递当前 `self.workspace`; 对 `g_idx_sort_indices` 改用 `replace_parameter(..., prefer_copy=True)` 注册, 确保每次后处理复制到同一存储。
3. 同步修复 FP8/NVFP4/MXFP4 密集和 MoE 共 8 个准备函数 (`marlin_utils_fp8.py`、`marlin_utils_fp4.py`、`compressed_tensors_moe_wna16_marlin.py`): 传递已有 `workspace`。
4. 在 `test_reload.py` 新增三个测试: 验证 `act-order` 内核的 `workspace` 和排序索引地址稳定、验证密集准备函数 `workspace` 地址稳定、验证不兼容时异常抛出。
5. 后续提交强化了不兼容异常并确认注册排序索引不影响 `reload accounting`。

关键文件:

- `vllm/model_executor/layers/quantization/utils/marlin_utils.py` (模块 量化工具; 类别 `source`; 类型 `data-contract`; 符号 `marlin_make_workspace_new`): 核心工具函数 `marlin_make_workspace_new` 增加 `existing` 参数和兼容性检查, 是其他所有站点修复的基础。
- `vllm/model_executor/kernels/linear/mixed_precision/marlin.py` (模块 量化内核; 类别 `source`; 类型 `data-contract`; 符号 `process_weights_after_loading`): 核心修复:

workspace 复用和 g_idx_sort_indices 注册，确保 CUDA 图地址稳定。

- tests/model_executor/model_loader/test_reload.py (模块 测试; 类别 test; 类型 test-coverage; 符号 _stub_marlin_ops, _make_act_order_marlin_kernel, _load_marlin_checkpoint_format_weights, _random_g_idx) : 新增三个回归测试, 验证 Marlin 运行时张量地址稳定性和兼容性, 无需 GPU 即可运行。
- vllm/model_executor/layers/quantization/utils/marlin_utils_fp8.py (模块 FP8 工具; 类别 source; 类型 data-contract; 符号 prepare_fp8_layer_for_marlin, prepare_fp8_moe_layer_for_marlin, prepare_mxfp8_layer_for_marlin, prepare_mxfp8_moe_layer_for_marlin) : 同步修复 FP8 密集和 MoE 准备函数的 workspace 分配。
- vllm/model_executor/layers/quantization/utils/marlin_utils_fp4.py (模块 FP4 工具; 类别 source; 类型 data-contract; 符号 prepare_fp4_layer_for_marlin, prepare_moe_fp4_layer_for_marlin) : 同步修复 NVFP4/MXFP4 密集和 MoE 准备函数的 workspace 分配。
- vllm/model_executor/layers/quantization/compressed_tensors/compressed_tensors_moe/compressed_tensors_moe_wna16_marlin.py (模块 压缩张量; 类别 source; 类型 data-contract; 符号 process_weights_after_loading) : 同步修复 compressed-tensors WNA16 MoE 路径的 workspace 分配。

关键符号: marlin_make_workspace_new, MarlinLinearKernel.process_weights_after_loading, prepare_fp8_layer_for_marlin, prepare_fp8_moe_layer_for_marlin, prepare_mxfp8_layer_for_marlin, prepare_mxfp8_moe_layer_for_marlin, prepare_fp4_layer_for_marlin, prepare_moe_fp4_layer_for_marlin, test_marlin_post_load_preserves_runtime_tensor_addresses, test_marlin_prepare_layer_preserves_workspace_address, test_marlin_make_workspace_new_rejects_incompatible_existing

关键源码片段

vllm/model_executor/layers/quantization/utils/marlin_utils.py

核心工具函数 marlin_make_workspace_new 增加 existing 参数和兼容性检查, 是其他所有站点修复的基础。

```
def marlin_make_workspace_new(
    device: torch.device,
    max_blocks_per_sm: int = 1,
    existing: torch.Tensor | None = None,
) -> torch.Tensor:
    # 线程块数 = SM 数 × 每 SM 最大线程块数
    sms = num_compute_units(device.index)
    size = sms * max_blocks_per_sm
    # 复用现有存储: 零填充后返回原指针, 保持 CUDA 图地址有效
    if existing is not None:
        if (
            existing.device != device
```

```

        or existing.dtype != torch.int
        or existing.numel() != size
    ):
        raise ValueError(
            f"Existing Marlin workspace is incompatible "
            f"(device={existing.device}, dtype={existing.dtype}, "
            f"numel={existing.numel()}; expected device={device}, "
            f"dtype={torch.int}, numel={size}). Reload must reuse the "
            f"workspace storage captured by CUDA graphs."
        )
    return existing.zero_()
# 首次调用：分配新工作区
return torch.zeros(size, dtype=torch.int, device=device, requires_grad=False)

```

vllm/model_executor/kernels/linear/mixed_precision/marlin.py

核心修复：workspace 复用和 g_idx_sort_indices 注册，确保 CUDA 图地址稳定。

```

def process_weights_after_loading(self, layer: torch.nn.Module) -> None:
    device = getattr(layer, self.w_q_name).device
    c = self.config
    # ... 初始化代码 ...
    size_k, size_n = c.partition_weight_shape
    if c.has_g_idx:
        padded_n, padded_k = size_n, size_k
    else:
        padded_n, padded_k = marlin_padded_nk(size_n, size_k, c.group_size)

    # 分配或复用 Marlin 工作区，确保 reload 时 CUDA 图地址不变
    self.workspace = marlin_make_workspace_new(
        device, existing=getattr(self, "workspace", None)
    )

    # ... 确定 g_idx 属性名和零点属性名 ...

    if c.has_g_idx:
        # act-order 路径：计算排序索引并注册为参数
        g_idx, g_idx_sort_indices = marlin_sort_g_idx(
            getattr(layer, self.w_gidx_name)
        )
        self._transform_param(layer, self.w_gidx_name, lambda _: g_idx)
        # 注册排序索引，使得 reload 时复制到相同存储
        replace_parameter(
            layer, "g_idx_sort_indices", g_idx_sort_indices, prefer_copy=True
        )
    else:
        setattr(layer, self.w_gidx_name, marlin_make_empty_g_idx(device))
        layer.g_idx_sort_indices = marlin_make_empty_g_idx(device)

```

tests/model_executor/model_loader/test_reload.py

新增三个回归测试，验证 Marlin 运行时张量地址稳定性和兼容性，无需 GPU 即可运行。

```
def test_marlin_post_load_preserves_runtime_tensor_addresses(monkeypatch, dist_init):
    from vllm.model_executor.layers.quantization.utils import marlin_utils

    _stub_marlin_ops(monkeypatch)
    kernel = _make_act_order_marlin_kernel()

    generator = torch.Generator().manual_seed(0)
    first_g_idx = _random_g_idx(generator)
    second_g_idx = _random_g_idx(generator)

    layer = torch.nn.Module()
    _load_marlin_checkpoint_format_weights(layer, first_g_idx)
    kernel.process_weights_after_loading(layer)

    workspace_ptr = kernel.workspace.data_ptr()
    sort_indices_ptr = layer.g_idx_sort_indices.data_ptr()

    # 模拟 reload: 加载不同 act-order 的权重并重新处理
    _load_marlin_checkpoint_format_weights(layer, second_g_idx)
    kernel.process_weights_after_loading(layer)

    # 断言 workspace 地址不变且被零填充
    assert kernel.workspace.data_ptr() == workspace_ptr
    assert torch.all(kernel.workspace == 0)
    # 断言排序索引地址不变
    assert layer.g_idx_sort_indices.data_ptr() == sort_indices_ptr
    # 断言排序索引的值对应新的 g_idx
    expected = marlin_utils.marlin_sort_g_idx(second_g_idx)[1]
    assert torch.all(layer.g_idx_sort_indices == expected)
```

评论区精华

1. Marlin workspace 兼容性检查: aoshen02 建议在 `marlin_make_workspace_new` 中当传入的 `existing workspace` 不兼容时直接 `raise` 而非静默重分配，避免在一级深度重现原始 bug。作者 RyanClark2k 同意并实施异常抛出，并添加单元测试验证。
 2. 注册 `g_idx_sort_indices` 为 `Parameter` 的影响: Codex 和 aoshen02 担心注册后可能导致 `load_numel_total` 膨胀，推迟处理。RyanClark2k 通过代码分析确认 `record_metadata_for_reloading` 在 `process_weights_after_loading` 之前运行，该张量不会进入 `restore` 集合，不影响 `accounting`。结论为无需额外改动。
- Marlin workspace 兼容性检查: 应抛出异常而非静默重分配 (design): 接受 `raise` 方案，已在后续 `commit` 中实现。
 - 注册 `g_idx_sort_indices` 为 `Parameter` 是否会影响 `reload accounting (correctness)`: 确认不影响，无需额外改动。

风险与影响

- 风险：主要风险：如果 workspace 大小随 GPU 改变（如迁移不同 SM 数量的 GPU）但已有 workspace 大小不变，会触发 ValueError，导致加载失败。这是预期的 fail-fast 行为，避免静默数据损坏。影响面广（6 个文件），但每个修改模式相同且经过 GPU 验证，风险可控。测试仅 CPU 环境运行，未覆盖真实 CUDA 图场景，但 PR body 包含独立 GPU 验证脚本。
- 影响：影响所有使用 Marlin 内核的量化模型（GPTQ、AWQ、compressed-tensors W4A16/W4A8）在 RL 权重重载场景。用户无需配置变更，加载行为兼容；CUDA 图使能时输出正确性从 bug 变为正确。对首次加载无性能影响，重载时因复用存储减少一次分配，略有微优化。
- 风险标记：CUDA 图地址有效性依赖，多路径同步修改，兼容性断言阻止未知迁移

关联脉络

- PR #48251 [Bugfix][Attention] Preserve post-load tensors across weight reloads: 类似模式的修复，针对注意力运行时张量存储保持，本 PR 参考了其 replace_parameter 模式。
- PR #48312 [RFC] Weight Reload Correctness for RL: 定义了权重重载正确性分类法，本 PR 专门解决了 Marlin workspace 和 sort indices 行。