

PR #48355 完整报告

vllm-project/vllm

feat: extended EPLB support for Mistral Large 3 and additional MoE backends

合并时间: 2026-08-07 20:32

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/48355>

执行摘要

- 一句话: EPLB 支持量化 MoE 与多模态模型, 修复重排精度崩溃
- 推荐动作: 值得精读。核心价值在于两点设计决策: 一是派生 per-expert 量化状态必须注册为 Parameter 并由 quant config 零拷贝别名, 保证 EPLB 就地在位重排后内核始终看到一致数据; 二是激活量化 scale 在 EP 组内做 amax 全归约, 从根上消除 quant/dequant 在专家迁移后的失配。配套的契约测试 (重排后必须等价于加载 permuted checkpoint) 写法很有借鉴意义, 适合作为量化 MoE + 动态重排类功能的回归模板。

功能与动机

EPLB 在推理运行期会沿 dim 0 就地重排 RoutedExperts 的每个注册 Parameter (见 RoutedExperts.get_expert_weights)。NVFP4/FP8 这类量化方法在 process_weights_after_loading 中派生的 per-expert scale 若不注册为 Parameter 并由 quant config 别名同一存储, 内核就会读到重排后的过期值, 产生 'silent accuracy corruption, no crash': PR body 给出的数据是 CuteDSL EPLB 7 次重排后 gsm8k 从 0.870/0.860 崩到 0.010/0.000, FP8 per-tensor 从 0.760/0.755 崩到 0.410/0.360。此外多模态模型把 MoE 语言模型嵌套在 .language_model 下, 自身不实现 MixtureOfExperts, 导致这些模型完全无法注册 EPLB。PR 的目标是让 EPLB 重排 registered Parameter 的结果与直接加载 permuted checkpoint 不可区分。

实现拆解

1. 统一 MoE 解析入口: 在 vllm/model_executor/models/interfaces.py 新增 get_mixture_of_experts_model(), 先判断自身是否 MixtureOfExperts, 再对 SupportsMultiModal 尝试 get_language_model() 递归判断; 同时将 MixtureOfExperts.moe_layers 类型从 Iterable 收紧为 Sequence (协变, 解决 mypy 下游报错)。vllm/v1/worker/gpu/eplb_utils.py 删除 _unwrap_moe, maybe_register_model、maybe_register_speculator、setup_from_mapping 全部改用新函数; vllm/v1/worker/gpu_model_runner.py 的 load_model() 对 drafter 与主模型统一收敛到该入口。
2. NVFP4 派生 scale 注册与 CuteDSL 视图化: compressed_tensors_moe_w4a4_nvfp4.py 用 replace_parameter 把 w13_weight_scale_2、w2_weight_scale_2 从普通 attribute 升级为注册 Parameter, 并新增 supports_eplb 白名单属性 (FLASHINFER_CUTEDSL / CUTEDSL_BATCHED / TRTLLM)。flashinfer_fp4_moe.py 的

`prepare_nvfp4_moe_layer_for_flashinfer_cutedsl` 只保留 `swizzle`, 不再内联做 MMA layout 转换; 新函数 `nvfp4_swizzled_scale_to_cutedsl_mma_view()` 在 `oracle/nvfp4.py::make_nvfp4_moe_quant_config` 构造时基于 `swizzled` 存储派生零拷贝视图, 并用 `data_ptr` 断言守护视图语义。

3. FP8 per-tensor alpha 注册: `oracle/fp8.py::make_fp8_moe_quant_config` 计算 `g1_alphas=w1_scale*a1_scale` 后注册为 layer 的 Parameter (`g1_alphas/g2_alphas`), `quant config` 直接引用同一存储, 并去掉多余的 `squeeze()`。
4. 激活量化 `scale` 的 EP 一致性: `quant_utils.py` 新增 `amax_for_moe_activation_quant(a_scale, enable_eplb)`, EPLB 开启时对 EP group 做 MAX all_reduce, 保证每个 rank 的量化 `scale` 与折入 dequant `alphas` 的值一致; `fp8_utils.process_fp8_input_tensor_strategy_moe` 增加 `enable_eplb` 参数并调用, native fp8 (`quantization/fp8.py`)、ModelOpt、compressed-tensors W8A8 FP8 三个调用点同步透传; NVFP4 两条 `prepare` 路径也替换为同一归约。
5. 测试配套: 新增 `tests/distributed/test_eplb_quant_scale_consistency.py` 契约测试 (NVFP4 双后端 + FP8 per-tensor, 含单 rank 分布式初始化以执行 EP all_reduce); 更新 CuteDSL NVFP4 内核测试、EPLB 运行器测试、TRTLLM hidden dim padding 测试以适配新签名。

关键文件:

- `tests/distributed/test_eplb_quant_scale_consistency.py` (模块 契约测试; 类别 test; 类型 test-coverage; 符号 `_RoutedExpertsStub`, `init`, `_expert_routing_tables`, `_make_moe_config`): 新增契约测试: 模拟 EPLB 就地重排所有 registered Parameter, 断言与直接加载 permuted checkpoint 完全一致, 并校验 `quant config` 与 Parameter 的 `data_ptr` 别名关系; 是本次修复的核心回归保障。
- `vllm/model_executor/layers/quantization/utils/flashinfer_fp4_moe.py` (模块 量化工具; 类别 source; 类型 data-contract; 符号 `nvfp4_swizzled_scale_to_cutedsl_mma_view`): NVFP4 CuteDSL 路径核心改动: MMA layout 转换后置为视图派生, 激活 `scale` 改为 EP 组 `amax` 归约。
- `vllm/model_executor/models/interfaces.py` (模块 模型接口; 类别 source; 类型 data-contract; 符号 `get_mixture_of_experts_model`): 新增 `get_mixture_of_experts_model()` 解析入口, 支持多模态模型中嵌套 MoE 语言模型的识别, 是 `eplb_utils` 与 `model_runner` 收敛的基础。
- `vllm/model_executor/layers/quantization/utils/quant_utils.py` (模块 量化工具; 类别 source; 类型 data-contract; 符号 `amax_for_moe_activation_quant`): 新增 `amax_for_moe_activation_quant()`: EPLB 开启时对 EP 组做 MAX all_reduce, 确保激活量化 `scale` 与融合进 dequant `alphas` 的数值在所有可能落脚的 rank 上一致。
- `vllm/v1/worker/gpu/eplb_utils.py` (模块 EPLB 控制; 类别 source; 类型 core-logic; 符号 `_unwrap_moe`): EPLBController 的模型 / 投机器注册与 `setup_from_mapping` 全部改用 `get_mixture_of_experts_model`, 使 VLM 嵌套 MoE 可被注册。
- `vllm/model_executor/layers/quantization/compressed_tensors/compressed_tensors_moe/compressed_tensors_moe_w4a4_nvfp4.py` (模块 NVFP4 量化; 类别 source; 类型 data-contract; 符号 `supports_eplb`): NVFP4 量化方法将 `w13/w2_weight_scale_2` 注册

为 Parameter，并新增 supports_eplb 属性限定已验证支持 EPLB 的后端白名单。

- vllm/v1/worker/gpu_model_runner.py (模块 GPU 运行器；类别 source；类型 data-contract) : load_model 中 drafter 与主模型的 EPLB 注册统一改用 get_mixture_of_experts_model，支持 VLM 嵌套 MoE 与嵌套 Drafter。
- vllm/model_executor/layers/fused_moe/oracle/fp8.py (模块 FP8 后端；类别 source；类型 data-contract) : FP8 per-tensor 的 g1/g2_alphas 从 detached 拷贝改为注册 Parameter，由 quant config 别名同一存储，EPLB 重排后内核可见更新。
- vllm/model_executor/layers/fused_moe/oracle/nvfp4.py (模块 NVFP4 后端；类别 source；类型 data-contract) : make_nvfp4_moe_quant_config 为 CuteDSL 后端构造 MMA scale 视图，替代 prepare 阶段的内联转换。
- vllm/model_executor/layers/quantization/utils/fp8_utils.py (模块 FP8 工具；类别 source；类型 data-contract) : process_fp8_input_tensor_strategy_moe 增加 enable_eplb 参数并调用 amax_for_moe_activation_quant，FP8 per-tensor 激活 scale 在 EPLB 下 EP 归约。
- vllm/model_executor/layers/quantization/compressed_tensors/compressed_tensors_moe/compressed_tensors_moe_w8a8_fp8.py (模块 FP8 量化；类别 source；类型 data-contract) : compressed-tensors W8A8 FP8 路径同步传递 enable_eplb 到 input scale 处理。
- vllm/model_executor/layers/quantization/fp8.py (模块 FP8 量化；类别 source；类型 data-contract) : native FP8 路径同步传递 enable_eplb，覆盖原生 fp8 量化 MoE。
- vllm/model_executor/layers/quantization/modelopt.py (模块 模型量化；类别 source；类型 data-contract) : ModelOpt FP8 路径同步传递 enable_eplb，覆盖 ModelOpt 配方。
- tests/kernels/moe/test_flashinfer_cutedsl_nvfp4_moe.py (模块 内核测试；类别 test；类型 test-coverage) : 适配 CuteDSL 路径的 prepare 函数签名变化与 MMA 视图后置。
- tests/v1/worker/test_gpu_model_runner_v2_eplb.py (模块 运行器测试；类别 test；类型 test-coverage) : 修复因 eplb_utils 删除 is_mixture_of_experts 导致的 AMD CI 失败，并适配新注册接口。
- tests/quantization/test_trtllm_nvfp4_hidden_dim_padding.py (模块 填充测试；类别 test；类型 test-coverage) : 覆盖 TRTLLM NVFP4 hidden dim padding 路径在本次改动后的回归。

关键符号: get_mixture_of_experts_model, amax_for_moe_activation_quant, nvfp4_swizzled_scale_to_cutedsl_mma_view, supports_eplb, make_fp8_moe_quant_config, process_fp8_input_tensor_strategy_moe, maybe_register_model, maybe_register_speculator, setup_from_mapping

关键源码片段

tests/distributed/test_eplb_quant_scale_consistency.py

新增契约测试: 模拟 EPLB 就地重排所有 registered Parameter，断言与直接加载 permuted checkpoint 完全一致，并校验 quant config 与 Parameter 的 data_ptr 别名关系；是本次修复的核心回归保障。

```
# tests/distributed/test_eplb_quant_scale_consistency.py
```

```
"""EPLB 重排一致性契约测试（节选）。
```

EPLB 通过 `RoutedExperts.get_expert_weights` 沿 dim 0 就地重排所有注册的 `Parameter`。量化方法必须把派生的 per-expert 张量注册为 `Parameter` 并让 `FusedMoEQuantConfig` 别名同一存储，内核才能看到重排后的值。

```
"""
```

```
def _simulate_eplb_rearrangement(layer: torch.nn.Module, perm: torch.Tensor) -> None:
```

```
    """按 EPLB 的方式，把每个注册 Parameter 沿 dim 0 就地重排。"""
```

```
    with torch.no_grad():
```

```
        for _, param in layer.named_parameters():
```

```
            param.copy_(param[perm])
```

```
def test_nvfp4_eplb_rearrangement_matches_reload(backend: str) -> None:
```

```
    # NVFP4 CuteDSL/TRTLLM MoE 后端需要 Blackwell (SM100)
```

```
    if not (current_platform.is_cuda()
```

```
            and current_platform.is_device_capability_family(100)):
```

```
        pytest.skip("NVFP4 CuteDSL/TRTLLM 后端需 Blackwell (SM100)")
```

```
    device = torch.device("cuda:0")
```

```
    perm = torch.tensor(EXPERT_PERMUTATION, device=device)
```

```
    raw = _make_raw_weights(device, generator)
```

```
    # 对照组：直接加载 permuted 专家顺序的 checkpoint
```

```
    raw_permuted = {name: value[perm].contiguous() for name, value in raw.items()}
```

```
    method, layer = _build_processed_layer(backend, raw, device)
```

```
    ref_method, ref_layer = _build_processed_layer(backend, raw_permuted, device)
```

```
    # EPLB 契约一：registered Parameter 必须 expert-major 且连续，
```

```
    # get_expert_weights 才能按 (E, -1) 视图就地重排专家切片
```

```
    for name, param in layer.named_parameters():
```

```
        assert param.is_contiguous(), f"{name} 不连续"
```

```
        assert param.shape[0] == NUM_EXPERTS, f"{name} 不是 expert-major"
```

```
    # EPLB 契约二：派生的 per-expert scale 必须由 quant config 别名
```

```
    # registered Parameter，任何脱离 Parameter 的拷贝都会在重排后过期
```

```
    quant_config = method.moe_quant_config
```

```
    params = dict(layer.named_parameters())
```

```
    assert quant_config.g1_alphas.data_ptr() == params["w13_weight_scale_2"].data_ptr()
```

```
    assert quant_config.g2_alphas.data_ptr() == params["w2_weight_scale_2"].data_ptr()
```

```
    assert quant_config.w1_scale.data_ptr() == params["w13_weight_scale"].data_ptr()
```

```
    assert quant_config.w2_scale.data_ptr() == params["w2_weight_scale"].data_ptr()
```

```
    # 就地重排后，所有状态必须等价于直接加载 permuted checkpoint
```

```
    _simulate_eplb_rearrangement(layer, perm)
```

```
    ref_params = dict(ref_layer.named_parameters())
```

```

for name in params:
    _assert_tensors_equal(params[name], ref_params[name], name)

# 内核可见的 quant config 张量 (CuteDSL MMA 视图别名同一存储) 也必须一致
ref_quant_config = ref_method.moe_quant_config
for name in QUANT_CONFIG_TENSORS:
    actual = getattr(quant_config, name)
    expected = getattr(ref_quant_config, name)
    if actual is not None:
        _assert_tensors_equal(actual, expected, f"quant_config.{name}")

```

vllm/model_executor/layers/quantization/utils/flashinfer_fp4_moe.py

NVFP4 CuteDSL 路径核心改动: MMA layout 转换后置为视图派生, 激活 scale 改为 EP 组 amax 归约。

```

# vllm/model_executor/layers/quantization/utils/flashinfer_fp4_moe.py

def nvfp4_swizzled_scale_to_cutedsl_mma_view(scale: torch.Tensor) -> torch.Tensor:
    """把 swizzled (E, M_padded, K_sf_padded) block-scale 张量以
    CuteDSL MoE 内核所需的 MMA layout 视图返回。

    返回值别名 scale 的存储: 注册 Parameter 的就地更新
    (权重重载、EPLB 重排) 对内核直接可见, 无需额外簿记。
    """
    from flashinfer.cute_dsl.utils import convert_sf_to_mma_layout

    num_experts, m_padded, k_sf_padded = scale.shape
    mma_view = convert_sf_to_mma_layout(
        scale.reshape(num_experts * m_padded, k_sf_padded),
        m=m_padded,
        k=k_sf_padded * 16,
        num_groups=num_experts,
        sf_vec_size=16,
    )
    # 若 flashinfer 实现变化导致不再返回视图, quant config 会在权重更新后过期
    assert mma_view.data_ptr() == scale.data_ptr(), (
        "convert_sf_to_mma_layout no longer returns a view of its input; "
        "the quant config would go stale after weight updates."
    )
    return mma_view

def prepare_nvfp4_moe_layer_for_flashinfer_cutedsl(
    layer: "RoutedExperts",
    w13: torch.Tensor,
    w13_scale: torch.Tensor,
    w13_scale_2: torch.Tensor,
    a13_scale: torch.Tensor,
    w2: torch.Tensor,

```

```

w2_scale: torch.Tensor,
w2_scale_2: torch.Tensor,
a2_scale: torch.Tensor,
) -> tuple(torch.Tensor, ...):
    # 前面: w13/w2 的 gate/up 行交错等 kernel 格式处理 (略)
    num_experts = w13.shape[0]
    enable_eplb = layer.moe_config.moe_parallel_config.enable_eplb
    # EPLB 下按 EP group 全归约 amax, 保证融合进 dequant alphas 的
    # 激活 scale 在每个专家落脚的 rank 上都与量化 scale 一致
    a13_scale = amax_for_moe_activation_quant(a13_scale, enable_eplb).repeat(
        num_experts
    )
    a2_scale = amax_for_moe_activation_quant(a2_scale, enable_eplb).repeat(num_experts)

    if layer.activation.is_gated:
        w13, w13_scale = reorder_w13_to_w31_for_flashinfer_cutedsl(
            layer.activation, w13, w13_scale
        )
        # Interleave up/gate rows for w13 weights and scales
        w13 = interleave_linear_and_gate(w13, group_size=64, dim=1)
        w13_scale = interleave_linear_and_gate(w13_scale, group_size=64, dim=1)

    # 只做 swizzle, 保留 expert-major 可注册布局; MMA layout 视图由
    # nvfp4_swizzled_scale_to_cutedsl_mma_view 在 quant config 构造时零拷贝派生
    w13_scale = swizzle_blockscale(w13_scale)
    w2_scale = swizzle_blockscale(w2_scale)

    return (
        w13, w13_scale, w13_scale_2, a13_scale,
        w2, w2_scale, w2_scale_2, a2_scale,
    )

```

评论区精华

Review 中 SageMoore 与作者 jdebache 围绕几处设计展开了交锋:

- fp8.py 的 squeeze 是否必要: SageMoore 问 'Why is the squeeze necessary here? What dimension are we trying to remove?', jdebache 回应这是从旧实现继承来的, 实际两个 scale 都已是 (num_experts,) 形状, squeeze 多余会移除。
- enable_eplb 传参 vs 读全局 config: SageMoore 建议直接在 amax_for_moe_activation_quant 内部读 get_current_vllm_config().parallel_config.enable_eplb, jdebache 倾向保持工具函数自包含、显式传参, SageMoore 最终认可 'I wouldn't do both, but I think the current form is fine.'
- CuteDSL MMA 视图转换是否仅 EPLB 启用: SageMoore 问 'Should we only be doing this when EPLB is enabled?', jdebache 解释转换本身不依赖 EPLB, 后置的目的是让 registered Parameter 保持 expert-major 可重排, 而 quant config 持零拷贝的 kernel 布局别名。

- 异常捕获过宽: SageMoore 认为 `get_language_model` 只需捕获 `NotImplementedError`, `AttributeError` 过宽; 最终实现收敛为只捕获 `NotImplementedError`。
- `moe_layers` 类型收紧: SageMoore 质疑 `Iterable` $\rightarrow$ `Sequence` 是否本 PR 所需, jdebache 说明是 `mypy` 协变问题, `Sequence` 协变才能通过类型检查。

此外 NickLucche 在 Issue 评论中报告了 AMD CI 失败: `vllm.v1.worker.gpu.eplb_utils` 已无 `is_mixture_of_experts` 属性, 属于 API 变更未同步到测试, 在最后一个 commit 'fixing tests' 中修复。

- `fp8.py` 中 `g1/g2_alphas` 的 `squeeze` 是否必要 (style): jdebache 确认两个 `scale` 来源已是 `(num_experts,)` 形状, `squeeze` 多余并移除。
- `amax_for_moe_activation_quant` 显式传参 vs 读全局 config (design): 保留显式 `enable_eplb` 参数, SageMoore 认可 'I wouldn't do both, but I think the current form is fine.'
- CuteDSL MMA 视图转换是否仅在 EPLB 下启用 (design): jdebache 解释转换本身非 EPLB 特有, 后置是为了让注册 `Parameter` 保持 `expert-major`, `quant config` 持零拷贝内核布局别名。
- `get_language_model` 异常捕获是否过宽 (correctness): 最终实现收敛为只捕获 `NotImplementedError`。
- `moe_layers` 从 `Iterable` 收紧为 `Sequence` 是否必要 (style): 保留该类型收紧。
- AMD CI 失败: `eplb_utils.is_mixture_of_experts` 属性不存在 (testing): jdebache 在 'fixing tests' commit 中同步更新测试引用。

风险与影响

- 风险:
 1. EP group 初始化约束: `amax_for_moe_activation_quant` 在 `enable_eplb` 时对 `get_ep_group().device_group` 做 `all_reduce`, 要求权重处理时分布式环境已初始化; 若 EP 组未建立会直接报错。新增测试必须 `_ensure_world1_distributed()` 建单 rank 环境才能运行。
 2. 依赖 flashinfer 视图语义: `nvfp4_swizzled_scale_to_cutedsl_mma_view` 用 `data_ptr` 断言 `convert_sf_to_mma_layout` 必须返回输入视图; flashinfer 升级后若改为拷贝, 断言会失败并暴露问题, 但属于硬性行为变更, 需要 CI 兜底。
 3. 非 EPLB 用户也受影响: CuteDSL 的 MMA 转换从 `prepare` 阶段后置到 `quant config` 构造阶段, 即使不启用 EPLB 也改变了执行路径与张量布局时序, 存在回归面; PR 中基线数据 (CuteDSL no EPLB 0.870/0.860 不变) 缓解了该担忧。
 4. 注册 `Parameter` 的连锁影响: `g1_alphas/g2_alphas` 与 `w13/w2_weight_scale_2` 变为 `registered Parameter` 后, 会进入 `named_parameters()`、`state_dict` 与优化器视角, 需确认权重加载、序列化路径与旧 `checkpoint` 兼容。
 5. API 变更同步风险: `eplb_utils` 删除 `is_mixture_of_experts` 引用后 AMD CI 即失败, 说明该模块被多路径引用; 后续第三方扩展若直接调用旧符号会破坏。- 影响: 影响面: EPLB + 量化 MoE 的所有用户, 尤其是 Mistral Large 3、RedHatAI/Qwen3-30B-A3B-NVFP4 (NVFP4 TRTLLM + CuteDSL) 和

DeepSeek-Coder-V2-Lite-Instruct-FP8 (FP8 per-tensor) 等模型组合。此前这些组合在 EPLB 发生重排后会静默产出错误结果，本 PR 将其修复为与 checkpoint 重排等价，属于生产质量级的正确性修复。同时多模态嵌套 MoE 模型（如 KimiK25 这类 wrapper）首次可注册 EPLB，扩展了功能覆盖。对团队而言，确立了 ' 量化派生状态必须注册为 Parameter 并由 quant config 别名 ' 的契约，后续新增量化后端时需遵循。改动集中在模型加载 / 权重处理路径，不影响推理热路径。 - 风险标记：核心路径变更，静默精度损坏修复，依赖 flashinfer 视图语义，EP 组初始化依赖，API 变更需同步测试

关联脉络

- PR #47106 [Kernel] Support Nvfp4 Cutedsl Moe Swiglu-oai and Relu2(non-gated)
Activation: 同一 CuteDSL NVFP4 MoE 内核 / 量化链路：本 PR 修改了 `prepare_nvfp4_moe_layer_for_flashinfer_cutedsl` 与 `test_flashinfer_cutedsl_nvfp4_moe.py`，并新增 MMA scale 视图函数，是对该 PR 内核支持的量化侧接力。
- PR #49764 [Quantization] Share online weight scales across TP: 同一 `quant_utils.py` 文件，同样解决量化 scale 跨 rank 一致性问题（该 PR 是 TP 归约，本 PR 是 EP 组 amax 归约），属于同一设计脉络。