

PR #52331 完整报告

vllm-project/vllm

[Test][LoRA] Speed up the LoRA test job

合并时间: 2026-08-15 01:08

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/52331>

执行摘要

- 一句话: LoRA 测试提速 3.7x: 引用上设备、FP8 向量化、跳过冗余清理
- 推荐动作: 值得精读, 尤其适合负责测试基建、CI 优化和 LoRA 模块的工程师。值得关注的设计决策包括: 按平台分流参考实现 (CPU 仅作为 XPU 的规避手段而非默认)、用零填充 + reshape + amax 把 blockwise 量化整体向量化、以及用 skip_global_cleanup 而非无操作 fixture 覆盖来跳过清理。建议合入后观察一次 XPU CI 结果, 并考虑把 quantize_to_fp8_blockwise 的写法沉淀为共享工具。

功能与动机

PR 正文明确指出: LoRA 任务是 CI 中最慢的任务之一, 且大部分时间花在非被测工作上 ("the LoRA job is one of the slower CI jobs and most of its time goes into work that is not under test")。典型例子是 test_punica_ops_fp8.py 单文件在 MI300X 上需要 768 秒, 其中大量时间消耗在逐 token、逐 block 的 Python 循环、每次用例后的全局清理, 以及参考实现与设备之间的 CPU 回拷上。本 PR 的目标就是移除这三类与测试断言无关的成本, 并明确承诺 "no test is dropped or weakened" (不丢弃、不削弱任何测试)。

实现拆解

1. 参考实现上设备 (tests/lora/test_punica_ops.py): 模块顶层新增 _REF_ON_CPU = current_platform.is_xpu() 与 _to_ref_device(), sgmv_shrink_for_nslices / sgmv_expand_for_nslices 不再无条件 .cpu(), 而是按平台把中间张量放到设备上直接做 per-LoRA matmul 循环; 只有 XPU 保留 CPU 计算并在最后 out_tensor.copy_(out) 回拷, 以规避 oneDNN/oneMKL 在大量 Triton 启动后返回错误结果的 XPU 专有问题。前序 PR #47534 引入的 per-LoRA 循环本身已解决约 8 GB 大张量物化问题, 因此 CPU 不再必要。该项使 4273 个测试从 475.74s 降至 37.46s (12.7x)。
2. FP8 blockwise 参考向量化 (tests/lora/test_punica_ops_fp8.py): quantize_to_fp8_blockwise、dequantize_fp8_blockwise 以及 generate_fp8_expand_data 的共享 scale 路径, 从多层嵌套 Python 循环改写为 F.pad 零填充 + reshape + amax 归约的整张量操作。核心不变量是 "零填充不改变 block 的 amax", 因此填充列在量化 / 反量化后切片丢弃, 结果与循环版本严格一致。该模式与 tests/kernels/quant_utils.py::native_per_token_group_quant_fp8 相同; 不直接复用的原因是后者要求最后一维可被 group size 整除, 且无法表达 (group_n, group_k) 的二维分块。该项使 1944 个测试从 768.40s 降至 15.50s (49.6x), 其中向量化贡献约 441s。

3. 跳过冗余的全局清理 (test_layers.py、test_fused_moe_lora_kernel.py、test_punica_ops_fp8.py) : 三个模块分别加上 `pytest.mark.skip_global_cleanup`。
test_layers.py 原本已有 `pytestmark = pytest.mark.skipif(...)`, 因此改为 mark 列表以同时保留 skipif; test_fused_moe_lora_kernel.py 的 TP 场景只在 `torch.multiprocessing.spawn` 的子进程内构建分布式环境, 父进程无可清理对象; FP8 punica 测试既不使用分布式也不使用 Ray。该标记在 tests/lora/conftest.py 中已有文档, 且 tests/lora/test_layers_utils.py 已在使用。清理项单独贡献约 311s。
4. 测试与验证配套: 未新增或删除任何测试用例; PR 给出逐文件 before/after 表格 (1944 / 4273 / 417 / 88 passed + 72 skipped 全部不变), 并对 test_punica_ops_fp8.py 的两项独立改动分别计时。维护者通过 /ci run 触发 Buildkite CI #83942 完成验证后合入。

关键文件:

- tests/lora/test_punica_ops.py (模块 测试套件; 类别 test; 类型 test-coverage; 符号 `_REF_ON_CPU`, `_to_ref_device`, `_bgmv_shrink`, `_bgmv_expand`) : 最核心的改动文件: 参考实现从强制 CPU 改为按平台上设备, 是 12.7x 提速的来源, 并以 `_REF_ON_CPU` 保留了 XPU 的 CPU 规避路径。
- tests/lora/test_punica_ops_fp8.py (模块 测试套件; 类别 test; 类型 test-coverage; 符号 `quantize_to_fp8_blockwise`, `dequantize_fp8_blockwise`, `generate_fp8_expand_data`) : 提速最显著的文件 (49.6x) : 把 blockwise FP8 量化 / 反量化与共享 scale 路径从逐块 Python 循环改写为零填充 + reshape + amax 的整张量操作, 并追加 `skip_global_cleanup`。
- tests/lora/test_layers.py (模块 测试套件; 类别 test; 类型 test-coverage; 符号 `pytestmark`) : 通过 `pytestmark` 列表保留 skipif 的同时追加 `skip_global_cleanup`, 展示已有 mark 与新增 mark 的组合方式, 测试时间从 81.11s 降至 48.49s。
- tests/lora/test_fused_moe_lora_kernel.py (模块 测试套件; 类别 test; 类型 test-coverage; 符号 `pytestmark`) : 新增 `skip_global_cleanup`: TP 用例的分布式环境只在 `multiprocessing` 子进程中构建, 父进程无可清理对象; 时间从 477.32s 降至 381.98s。

关键符号: `_to_ref_device`, `_bgmv_shrink`, `_bgmv_expand`, `sgmv_shrink_for_nslices`, `sgmv_expand_for_nslices`, `quantize_to_fp8_blockwise`, `dequantize_fp8_blockwise`, `generate_fp8_expand_data`

关键源码片段

tests/lora/test_punica_ops.py

最核心的改动文件: 参考实现从强制 CPU 改为按平台上设备, 是 12.7x 提速的来源, 并以 `_REF_ON_CPU` 保留了 XPU 的 CPU 规避路径。

```
# tests/lora/test_punica_ops.py (本次合并后版本, 节选)
# 除 XPU 外, 参考 matmul 直接在设备上计算, 避免每轮测试的 CPU 回拷。
_REF_ON_CPU = current_platform.is_xpu()
```

```
def _to_ref_device(tensor: torch.Tensor) -> torch.Tensor:
    # XPU 保留 CPU 路径: oneDNN 与 oneMKL 在大量 Triton 启动后可能算错,
    # 这是 XPU 特有的数值问题; 其余平台直接返回原张量。
```

```
return tensor.cpu() if _REF_ON_CPU else tensor
```

```
def _bgmv_shrink(inputs, lora_weight, output, seq_len_tensor, lora_indices, scaling=1.0):
```

```
    """内存友好的 shrink 参考：按 LoRA 分组的 matmul 循环。
```

```
    output[mask] = scaling * inputs[mask] @ weight.T
```

```
    前序改动 #47534 引入的 per-LoRA 循环避免了物化约 8 GB 的
```

```
    lora_weight[exploded_indices], 因此不再依赖 CPU 来省内存。
```

```
    """
```

```
    exploded = torch.repeat_interleave(lora_indices, seq_len_tensor)
```

```
    for lid in exploded.unique():
```

```
        if lid < 0:
```

```
            continue
```

```
        mask = exploded == lid
```

```
        inp = inputs[mask].to(output.dtype)
```

```
        w = lora_weight[lid].to(output.dtype)
```

```
        output[mask] = scaling * (inp @ w.T)
```

```
def sgmv_shrink_for_nslices(nslices, inputs_tensor, lora_weights_lst, out_tensor,
```

```
                            b_seq_start_loc, seq_len_tensor, prompt_lora_mapping,
```

```
                            batches, max_seq_length, num_tokens, scaling):
```

```
    """参考实现：所有中间张量经 _to_ref_device 统一落到目标设备。"""
```

```
    inputs = _to_ref_device(inputs_tensor)
```

```
    seq_len = _to_ref_device(seq_len_tensor)
```

```
    mapping = _to_ref_device(prompt_lora_mapping)
```

```
    out = _to_ref_device(out_tensor)
```

```
    for index in range(nslices):
```

```
        _bgmv_shrink(inputs, _to_ref_device(lora_weights_lst[index]),
```

```
                    out[index], seq_len, mapping, scaling=scaling)
```

```
    # 仅 XPU 需要把结果拷回设备；其余平台直接在设备上写 out。
```

```
    if _REF_ON_CPU:
```

```
        out_tensor.copy_(out)
```

tests/lora/test_punica_ops_fp8.py

提速最显著的文件（49.6x）：把 blockwise FP8 量化 / 反量化与共享 scale 路径从逐块 Python 循环改写为零填充 + reshape + amax 的整张量操作，并追加 skip_global_cleanup。

```
# tests/lora/test_punica_ops_fp8.py（本次合并后版本，节选）
```

```
def quantize_to_fp8_blockwise(tensor, group_n, group_k):
```

```
    """整张量版本：零填充补齐到整块后，用 reshape + amax 一次算完所有 scale。
```

```
    零填充不会改变 block 的 amax（abs(0) = 0），因此量化后把填充列切片
```

```
    丢弃即可，结果与逐块循环完全一致。
```

```
    """
```

```
    if tensor.ndim == 2:
```

```
        M, K = tensor.shape
```

```
        n_blocks_k = math.ceil(K / group_k)
```

```

# 只补最后一维 K 到整块。
padded = F.pad(tensor.float(), (0, n_blocks_k * group_k - K))
blocks = padded.view(M, n_blocks_k, group_k)
scale = blocks.abs().amax(dim=-1).clamp(min=1e-12) / FP8_MAX
fp8_tensor = ((blocks / scale.unsqueeze(-1))
               .clamp(FP8_MIN, FP8_MAX).to(FP8_DTYPE))
return fp8_tensor.view(M, -1)[:K].contiguous(), scale
if tensor.ndim == 3:
    L, N, K = tensor.shape
    n_blocks_n = math.ceil(N / group_n)
    n_blocks_k = math.ceil(K / group_k)
    # N 与 K 两个维度都要补齐到整块。
    padded = F.pad(tensor.float(),
                    (0, n_blocks_k * group_k - K, 0, n_blocks_n * group_n - N))
    blocks = padded.view(L, n_blocks_n, group_n, n_blocks_k, group_k)
    # amax 在 (group_n, group_k) 两个块轴上同时归约。
    scale = blocks.abs().amax(dim=(2, 4)).clamp(min=1e-12) / FP8_MAX
    fp8_tensor = ((blocks / scale[:, :, None, :, None])
                  .clamp(FP8_MIN, FP8_MAX).to(FP8_DTYPE)
                  .view(L, n_blocks_n * group_n, n_blocks_k * group_k))
    return fp8_tensor[:, :N, :K].contiguous(), scale
raise ValueError(f"Unsupported tensor ndim: {tensor.ndim}")

```

```

def dequantize_fp8_blockwise(fp8_tensor, scale, group_n, group_k, output_dtype=torch.bfloat16)
:

```

```

"""与 quantize 对称：先零填充到整块，反量化后再裁掉填充列。"""
if fp8_tensor.ndim == 2:
    M, K = fp8_tensor.shape
    n_blocks_k = math.ceil(K / group_k)
    padded = F.pad(fp8_tensor.float(), (0, n_blocks_k * group_k - K))
    blocks = padded.view(M, n_blocks_k, group_k)
    out = (blocks * scale.float()).unsqueeze(-1).to(output_dtype)
    return out.view(M, -1)[:K].contiguous()
if fp8_tensor.ndim == 3:
    L, N, K = fp8_tensor.shape
    n_blocks_n = math.ceil(N / group_n)
    n_blocks_k = math.ceil(K / group_k)
    padded = F.pad(fp8_tensor.float(),
                    (0, n_blocks_k * group_k - K, 0, n_blocks_n * group_n - N))
    blocks = padded.view(L, n_blocks_n, group_n, n_blocks_k, group_k)
    out = (blocks * scale.float())[:, :, None, :, None].to(output_dtype)
    out = out.view(L, n_blocks_n * group_n, n_blocks_k * group_k)
    return out[:, :N, :K].contiguous()
raise ValueError(f"Unsupported tensor ndim: {fp8_tensor.ndim}")

```

评论区精华

本次 review 评论极少且无技术争议：claude[bot] 提示 fork 来源 PR 的自动审核被禁用（可手动 [/claude review](#)）；维护者 AndreasKaratzas 触发 [/ci run](#) 并通过 Buildkite CI #83942，最终以 APPROVED 和 "LGTM" 合入。设计权衡（XPU 保留 CPU 路径、零填充不改变 `amax`、不直接复用 `native_per_token_group_quant_fp8` 的原因）主要记录在 PR 正文中，未产生讨论分歧。

- 触发 CI 验证 (other): CI 验证通过，未暴露测试失败。
- 维护者审核与自动 review 禁用 (other): PR 获批并合并，无技术争议。

风险与影响

- 风险：

1. XPU 路径未实测：`_REF_ON_CPU = current_platform.is_xpu()` 是唯一平台开关，测速数据全部来自 MI300X；若平台判断与实际 `oneDNN/oneMKL` 行为不符，XPU 上可能走错路径。建议在 XPU CI 补跑 `tests/lora/test_punica_ops.py`。
2. 零填充等价性：向量化正确性依赖 "零填充不改变 `block` 的 `amax`" 这一不变量；当前只用于测试内部且测试数量不变，但未来若被复用到非零填充场景需重新验证。
3. 跳过全局清理：`skip_global_cleanup` 可能掩盖跨用例的分布式环境泄漏；三个文件要么由 `dist_init` 自行清理、要么从不构建分布式环境，且该标记已有先例（`tests/lora/test_layers_utils.py`），风险可控。
4. 影响面：仅测试代码，无运行时、API 或 schema 变更。- 影响：对 CI 流程影响直接且显著：LoRA 测试 job 在单台 MI300X 上从约 30 分钟降至约 8 分钟，节省约 22 分钟；由于 CI 主机共享、AMD job 只占八分之一 GPU，实际收益需 CI 实测确认。对开发者而言，本地执行 `tests/lora/` 四件套的反馈周期大幅缩短（如 `test_punica_ops_fp8.py` 从 768 秒降至 15 秒）。对系统运行时零影响；对团队而言，这是可复用的测试基建优化范式（平台分流、向量化参考实现、精准清理标记）。- 风险标记：XPU 未实测，零填充等价性依赖，清理跳过需回归，纯测试改动

关联脉络

- PR #47534（标题未在材料中提供，PR 正文引用）：PR 正文明确说明本 PR 第一项优化保留了 #47534 引入的 `per-LoRA matmul` 循环，只是把参考实现从 CPU 移到设备；它是本次改动的直接前置。
- PR #52237 [UT] fix device of `test_outputs.py`: 同为测试基建改进，将设备判断改为 `current_platform` 以支持非 CUDA 平台，与本 PR 的 XPU 平台分支思路一致。
- PR #52252 [CI] Increase extended generation test timeout: 同属 CI 耗时治理：一个放宽超时、一个压缩耗时，说明团队正在持续优化 CI 时间线。