

PR #41934 完整报告

vllm-project/vllm

[Hardware][XPU] Register batch-invariant kernels for XPU

合并时间: 2026-07-15 23:19

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/41934>

执行摘要

- 一句话: 为 XPU 注册批次不变核函数, 支持强化学习
- 推荐动作: 值得精读, 尤其是其跨平台 dispatch key 抽象和测试基础设施的统一设计。评审中关于 allow_override 的讨论提醒了 dispatch 冲突处理, 是值得关注的设计细节。建议在后续 PR 中尽快补齐 matmul/linear 的注册, 并移除 E2E 测试的 xfail。

功能与动机

批次不变性是强化学习 (RL) 在 vLLM 中工作的必要条件。XPU 用户需要该特性来支持 RL 流程。PR 是使 XPU 支持 RL 的第一步, 后续会逐步完善余下的 kernel 注册和性能优化。

实现拆解

1. 修改核心注册入口 `vllm/model_executor/layers/batch_invariant.py` 中的 `enable_batch_invariant_mode()`:
 - 提取 `current_platform.dispatch_key`, 替代原先硬编码的 "CUDA", 使注册可扩展至其他后端。
 - 添加 `elif current_platform.is_xpu():` 分支, 注册 mm 和 addmm 的批次不变 Triton kernel, 并固定 `_fp16_block_size_n = 128`。
 - 将公共覆盖 (softmax、log_softmax、mean、bmm) 从 "CUDA" 改为使用 key 统一分发。
 - 将 cuBLAS 精度相关设置包裹在 `if current_platform.is_cuda():` 下, 避免在 XPU 上误设。
2. 重构测试基础 `tests/v1/determinism/utils.py`:
 - 引入 DeviceConfig NamedTuple, 描述每个设备的可用性和支持的后端列表。
 - 定义 DEVICE_BACKENDS 字典, 键为 "cuda" 和 "xpu", 初始值包括各自的可用判断和支持后端。
 - 当使用 MLA 模型时, 同时影响 CUDA 和 XPU 的后端列表。
 - 新增 skip_if_not_cuda 标记, 用于仅 CUDA 适用的测试。
3. 调整现有测试:
 - `tests/v1/determinism/test_rms_norm_batch_invariant.py`: 将原 `@skip_unsupported` 改为 `@skip_if_not_cuda`, 避免在 XPU 上运行 (因为 XPU 上尚未有 CUDABased RMSNorm 对比)。新增 `test_rms_norm_batch_invariance` 函数, 直接验证同一行在不同批次大小时输出是否一致, 此测试对所有支持平台有效。

- tests/v1/determinism/test_batch_invariance.py: 在 test_v1_generation_is_deterministic_across_batch_sizes_with_needle 和 test_logprobs_bitwise_batch_invariance_bs1_vs_bsN 开头添加 XPU 检测, 若为 XPU 则 pytest.xfail。
- tests/v1/determinism/test_online_batch_invariance.py 和 tests/v1/determinism/test_nvfp4_batch_invariant.py: 将 @skip_unsupported 替换为 @skip_if_not_cuda。

4. 文档更新[docs/features/batch_invariance.md](#): 添加 XPU 实验性支持的章节, 说明注意后端选择。

关键文件:

- vllm/model_executor/layers/batch_invariant.py (模块 批次不变性; 类别 source; 类型 core-logic; 符号 enable_batch_invariant_mode, init_batch_invariance) : 核心修改: 添加 XPU 分支, 将 dispatch key 从硬编码的 'CUDA' 改为动态获取, 并为 XPU 注册 mm、addmm、softmax 等批次不变 kernel。
- tests/v1/determinism/utils.py (模块 测试工具; 类别 test; 类型 test-coverage; 符号 DeviceConfig, DEVICE_BACKENDS, BACKENDS, skip_unsupported) : 测试基础设施重构: 引入 DeviceConfig 统一管理设备可用性和后端列表, 新增 skip_if_not_cuda 标记, 为多平台测试提供基础。
- tests/v1/determinism/test_rms_norm_batch_invariant.py (模块 RMS Norm 测试; 类别 test; 类型 test-coverage; 符号 test_rms_norm_batch_invariance) : 新增 test_rms_norm_batch_invariance 测试, 直接验证批次不变性质; 同时将多个 CUDA 特定测试改为使用 @skip_if_not_cuda, 避免在 XPU 上运行。
- tests/v1/determinism/test_batch_invariance.py (模块 E2E 测试; 类别 test; 类型 test-coverage) : E2E 测试中添加 XPU xfail 处理, 避免因未注册完整 kernel 导致测试失败。
- tests/v1/determinism/test_online_batch_invariance.py (模块 在线测试; 类别 test; 类型 test-coverage) : 将 @skip_unsupported 替换为 @skip_if_not_cuda, 确保在线批次不变性测试仅在 CUDA 上运行。
- tests/v1/determinism/test_nvfp4_batch_invariant.py (模块 NVFP4 测试; 类别 test; 类型 test-coverage) : 将 @skip_unsupported 替换为 @skip_if_not_cuda, 确保 NVFP4 批次不变测试仅在 CUDA 上运行。

关键符号: enable_batch_invariant_mode, init_batch_invariance, test_rms_norm_batch_invariance

关键源码片段

[vllm/model_executor/layers/batch_invariant.py](#)

核心修改: 添加 XPU 分支, 将 dispatch key 从硬编码的 'CUDA' 改为动态获取, 并为 XPU 注册 mm、addmm、softmax 等批次不变 kernel。

```
def enable_batch_invariant_mode():
    global _batch_invariant_MODE, _batch_invariant_LIB
```

```

global _fp16_block_size_n

if _batch_invariant_MODE:
    return

_batch_invariant_MODE = True
_batch_invariant_LIB = torch.library.Library("aten", "IMPL")

# 使用 dispatch_key 替代硬编码 "CUDA", 便于扩展到 XPU 等后端
key = current_platform.dispatch_key

if current_platform.is_cuda():
    if current_platform.is_device_capability_family(80):
        # Ampere 需要 Triton persistent matmul 覆盖
        _batch_invariant_LIB.impl("aten::mm", mm_batch_invariant, key)
        _batch_invariant_LIB.impl("aten::addmm", addmm_batch_invariant, key)
        _batch_invariant_LIB.impl("aten::matmul", matmul_batch_invariant, key)
        _batch_invariant_LIB.impl("aten::linear", linear_batch_invariant, key)
    else:
        # Hopper/Blackwell 仅关闭 split-k
        os.environ["CUBLAS_WORKSPACE_CONFIG"] = ":16:8"
        os.environ["CUBLASLT_WORKSPACE_SIZE"] = "1"
        _fp16_block_size_n = 256 if get_max_shared_memory_bytes() > 106496 else 128
elif current_platform.is_xpu():
    # XPU 仅注册已验证的 mm 和 addmm, matmul/linear 待后续实现
    _batch_invariant_LIB.impl("aten::mm", mm_batch_invariant, key)
    _batch_invariant_LIB.impl("aten::addmm", addmm_batch_invariant, key)
    # TODO: register matmul and linear for XPU
    # once suitable Triton kernels are implemented
    _fp16_block_size_n = 128 # XPU 上使用保守值

# 以下覆盖对所有平台统一注册 (CUDA 和 XPU)
_batch_invariant_LIB.impl("aten::_log_softmax", _log_softmax_batch_invariant, key)
_batch_invariant_LIB.impl("aten::softmax", softmax_batch_invariant, key)
_batch_invariant_LIB.impl("aten::_softmax", softmax_batch_invariant, key)
_batch_invariant_LIB.impl("aten::mean.dim", mean_batch_invariant, key)
# bmm 需要 allow_override=True 以应对 Torch 内置注册
_batch_invariant_LIB.impl("aten::bmm", bmm_batch_invariant, key, allow_override=True)
torch.bmm = bmm_batch_invariant

reduced_precision_val = (
    (False, False) if is_torch_equal_or_newer("2.10.0") else False
)
# CUDA 特定的低精度归约设置
if current_platform.is_cuda():
    torch.backends.cuda.matmul.allow_fp16_reduced_precision_reduction = reduced_precision_val
    torch.backends.cuda.matmul.allow_bf16_reduced_precision_reduction = reduced_precision_val

```

```
torch.backends.cuda.preferred_blas_library(backend="cublaslt")
```

tests/v1/determinism/utils.py

测试基础设施重构：引入 DeviceConfig 统一管理设备可用性和后端列表，新增 skip_if_not_cuda 标记，为多平台测试提供基础。

```
class DeviceConfig(NamedTuple):
    """每个设备的可用性和支持后端列表"""
    available: bool
    backends: list[str]

# 为 CUDA 和 XPU 分别定义配置
DEVICE_BACKENDS: dict[str, DeviceConfig] = {
    "cuda": DeviceConfig(
        available=current_platform.is_cuda()
        and current_platform.has_device_capability(80), # 需要 Sm80+
        backends=["FLASH_ATTN", "TRITON_ATTN", "FLEX_ATTENTION"],
    ),
    "xpu": DeviceConfig(
        available=current_platform.is_xpu() and HAS_TRITON,
        backends=["TRITON_ATTN"], # XPU 当前仅支持 Triton Attention
    ),
}

# BACKENDS 通过集合去重后排序，仅包含实际可用设备的后端
BACKENDS: list[str] = sorted(
    {b for cfg in DEVICE_BACKENDS.values() if cfg.available for b in cfg.backends}
)

skip_unsupported = pytest.mark.skipif(
    not any(cfg.available for cfg in DEVICE_BACKENDS.values()),
    reason="Requires CUDA >= Ampere (SM80) or Intel XPU with Triton",
)

skip_if_not_cuda = pytest.mark.skipif(
    not DEVICE_BACKENDS["cuda"].available,
    reason="Requires CUDA >= Ampere (SM80)",
)
```

tests/v1/determinism/test_rms_norm_batch_invariant.py

新增 test_rms_norm_batch_invariance 测试，直接验证批次不变性质；同时将多个 CUDA 特定测试改为使用 @skip_if_not_cuda，避免在 XPU 上运行。

```
@skip_unsupported # 对所有支持平台运行
@pytest.mark.parametrize("dtype", [torch.float16, torch.bfloat16])
def test_rms_norm_batch_invariance(dtype):
    """验证同一行在不同批次邻居下 RMS Norm 输出一致：核心批次不变性质"""
    device = torch.device(DEVICE_TYPE)
    torch.manual_seed(42)
```

```

hidden_size = 2048
eps = 1e-6

weight = torch.randn(hidden_size, dtype=dtype, device=device)
row = torch.randn(1, hidden_size, dtype=dtype, device=device)

# 单独一行计算 RMS Norm
out_single = rms_norm_batch_invariant(row, weight, eps=eps)

# 将同一行嵌入更大批次（位置 4）
batch = torch.randn(8, hidden_size, dtype=dtype, device=device)
batch[4] = row[0]
out_batch = rms_norm_batch_invariant(batch, weight, eps=eps)

assert torch.equal(out_single[0], out_batch[4]), (
    "rms_norm output for a row differs when batch context changes"
)

```

评论区精华

1. `allow_override=True` 的必要性 (gemini-code-assist[bot])：指出 XPU 端 bmm 注册需要加 `allow_override=True`，否则会因重复注册抛出 `RuntimeError`。作者随后采纳。
 2. 是否新增独立测试文件 (yewentao256, xuechendi)：评审者建议不要创建 `test_xpu_batch_invariant.py`，而应复用现有测试。作者最终将测试合并到 `test_rms_norm_batch_invariant.py` 中，并重构了 skip 机制。
 3. E2E 测试标记为 xfail 的合理性 (yewentao256)：评审者认为“我们期望它是确定的”，但作者解释这是逐步启用的策略，先注册一部分 kernel，后续再补充。最终保留 xfail。
 4. 平台检查是否需要包含 ROCm (yewentao256)：作者最初的 `enable_batch_invariant_mode` 开头检查了 `is_cuda()` or `is_xpu()`，评审者询问是否考虑了 ROCm；作者删除该检查，让其他平台继续走原逻辑（无覆盖）。
- `allow_override=True` 的必要性 (correctness): 添加 `allow_override=True` 到 bmm 注册调用中。
 - 是否创建新的测试文件 (design): 删除独立测试文件，改为修改现有测试文件和重构 `utils.py`。
 - E2E 测试标记为 xfail 的合理性 (design): 保留 xfail，计划在后续 PR 中补齐 kernel 注册后移除。
 - 平台检查是否包含 ROCm (correctness): 移除平台限制，`enable_batch_invariant_mode` 对 ROCm 保持不变。

风险与影响

- 风险：
 - 不完整的批次不变性支持：未注册 `matmul` 和 `linear` 的 Triton kernel，导致 XPU 上批次不变性仅覆盖部分算子，E2E 测试无法通过。如果用户期望完全确定性，可能会误以为 XPU 批次不变性不可用。

- CUDA 兼容性回归：修改了 `enable_batch_invariant_mode` 中 `dispatch key` 的使用，将原先硬编码的 "CUDA" 改为 `key`。如果 `dispatch_key` 返回非预期值，可能导致 CUDA 上的覆盖注册失败。但根据代码，CUDA 分支保持了原有行为，风险较低。
- ROCm 未验证：代码中未对 ROCm 做特殊处理，但 ROCm 也不会进入 XPU 分支；如果 ROCm 用户调用 `enable_batch_invariant_mode`，会沿用原有的 CUDA 逻辑（因为 `is_cuda()` 可能为 `False`），但不会错误注册 XPU kernel。
- 测试基础设施重构影响：`utils.py` 中的设备后端配置被重构，可能影响其他测试文件（如 `test_batch_invariance.py` 等）。但这些改动主要是让测试可以感知 XPU，已有测试的行为通过 `skip_if_not_cuda` 保持不变。
- 影响：用户影响：对 Intel XPU 用户，批次不变性功能首次可用（尽管不完整），为 RL workflow 奠定基础。对 CUDA 用户，无行为变化。对 ROCm 等其他平台用户，无影响。系统影响：仅影响 `batch invariance` 模块及其测试。核心 `serving` 链路的其他部分未修改。团队影响：为 XPU 团队后续性能优化和完整 kernel 覆盖提供了基础结构和测试框架。
- 风险标记：不完整 kernel 注册，E2E 测试 `xfail`，ROCm 兼容性未验证，`dispatch key` 抽象风险

关联脉络

- 暂无明显关联 PR