

PR #49444 完整报告

vllm-project/vllm

[Misc] Enable test_silu_mul_fp8_quant_deep_gemm on XPU

合并时间: 2026-08-11 23:19

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/49444>

执行摘要

- 一句话: 修复 XPU 上 FP8 SiLU 内核崩溃并启用对应测试
- 推荐动作: 值得快速精读, 作为平台抽象 (current_platform) 用法的范例: "能力缺失时优雅回退而非断言失败"。对维护多后端 (CUDA/ROCm/XPU) 的团队有参考价值, 尤其是 has_device_capability() 替换 get_device_capability()+assert 的写法值得在同类平台分支中推广。

功能与动机

PR body 明确指出: persistent_masked_m_silu_mul_quant() 调用 current_platform.get_device_capability().to_int() 并断言其非 None, 而 XpuPlatform 返回 None, 导致该 wrapper 在 XPU 上直接崩溃而非回退到 Triton 路径。此次变更的目标是让 deep_gemm 的 SiLU + Mul + FP8 量化融合内核在 XPU 上可回退、可测试, 使 Intel GPU 用户获得与 ROCm 相同的 Triton 路径覆盖。

实现拆解

1. 平台判断重构 (源码): 在 vllm/model_executor/layers/fused_moe/experts/batched_deep_gemm_moe.py 的 persistent_masked_m_silu_mul_quant() 中, 删除 get_device_capability() + assert + to_int() 三段式判断, 改为 if current_platform.is_cuda() and current_platform.has_device_capability(80)。has_device_capability() 在平台无能力概念时返回 False, 因此 CUDA sm_80+ 仍走 C++ 原生内核, ROCm/XPU 自动落入 Triton 回退路径; 删除 assert 同时消除了 XPU 上的崩溃窗口。
2. 测试去除 CUDA 硬编码: tests/kernels/moe/test_silu_mul_fp8_quant_deep_gemm.py 引入 DEVICE = current_platform.device_type, 替换 pack_scales、ref_with_scale_fmt、token_random 以及 tokens_per_expert 分配中的 device="cuda", 参考实现张量跟随输入设备, 测试可在任意受支持平台上执行。
3. 平台门控与容差扩展: 新增模块级 pytestmark = pytest.mark.skipif(not (is_cuda_alike() or is_xpu()), ...) 整文件门控; UE8M0 打包 scale 用例仍仅在 is_cuda() 时追加 (依赖 C++ 内核); 由于 XPU 与 ROCm 共用 Triton 回退内核, 断言处以 is_rocm() or is_xpu() 统一放宽为 atol=32.0 / rtol=0.2 的 1-ULP FP8 容差。
4. 验证: 作者在 Intel Arc Pro B70 与 H200 上跑通 23/23 用例; jikunshang 触发 Buildkite CI #83358 全量验证。无配置、schema 或部署配套改动。

关键文件：

- vllm/model_executor/layers/fused_moe/experts/batched_deep_gemm_moe.py (模块 MoE 内核；类别 source；类型 core-logic；符号 persistent_masked_m_silu_mul_quant)：核心修复文件：persistent_masked_m_silu_mul_quant() 的平台判断从 get_device_capability()+assert 改为 is_cuda() and has_device_capability(80)，消除 XPU 崩溃并让 ROCm/XPU 统一走 Triton 回退路径。
- tests/kernels/moe/test_silu_mul_fp8_quant_deep_gemm.py (模块 单元测试；类别 test；类型 test-coverage；符号 test_silu_mul_fp8_quant_deep_gemm, pack_scales, ref_with_scale_fmt, token_random)：测试平台化：引入 DEVICE 常量替换硬编码 "cuda"，新增 CUDA-alike/XPU 的 skipif 门控，并将 1-ULP FP8 容差扩展到 XPU，使测试可在 Intel GPU 上运行并通过。

关键符号：persistent_masked_m_silu_mul_quant, test_silu_mul_fp8_quant_deep_gemm, pack_scales, ref_with_scale_fmt, token_random

关键源码片段

vllm/model_executor/layers/fused_moe/experts/batched_deep_gemm_moe.py

核心修复文件：persistent_masked_m_silu_mul_quant() 的平台判断从 get_device_capability()+assert 改为 is_cuda() and has_device_capability(80)，消除 XPU 崩溃并让 ROCm/XPU 统一走 Triton 回退路径。

```
def persistent_masked_m_silu_mul_quant(
    y: torch.Tensor,
    tokens_per_expert: torch.Tensor,
    group_size: int = 128,
    quant_scale_fmt: DeepGemmQuantScaleFMT = DeepGemmQuantScaleFMT.FLOAT32,
) -> tuple[torch.Tensor, torch.Tensor]:
    """SiLU + Mul + FP8 量化融合内核的统一入口 (deep_gemm 通道) 。

    CUDA sm_80+ 走原生 C++ 内核，ROCM / XPU 走 Triton 回退。
    """
    E, T, H2 = y.shape # (E, T, 2*H): gate 与 up 拼接后的输入
    H = H2 // 2
    G = (H + group_size - 1) // group_size

    tokens_per_expert = tokens_per_expert.to(device=y.device, dtype=torch.int32)
    fp8_dtype = current_platform.fp8_dtype()
    y_q = torch.empty((E, T, H), dtype=fp8_dtype, device=y.device)
    # ... 中间省略 y_s (量化 scale) 按 quant_scale_fmt 布局分配并清零 ...
    ceil_ue8m0 = quant_scale_fmt in (
        DeepGemmQuantScaleFMT.FLOAT32_CEIL_UE8M0,
        DeepGemmQuantScaleFMT.UE8M0,
    )

    # 关键变更：旧代码先取 get_device_capability() 再断言非 None。
```

```

# XpuPlatform 没有 " 设备计算能力 " 概念, 返回 None 直接触发断言崩溃。
# has_device_capability(80) 在平台无能力概念时返回 False,
# 因此 CUDA sm_80+ 走 C++ 内核, ROCm / XPU 自动落入 Triton 路径。
if current_platform.is_cuda() and current_platform.has_device_capability(80):
    torch.ops._C.persistent_masked_m_silu_mul_quant(
        y, tokens_per_expert, y_q, y_s, ceil_ue8m0
    )
else:
    # Triton 回退 (ROCm 与 XPU 共用) : C++ 内核受
    # activation_kernels.cu 的 #ifndef USE_ROCM 保护, 且未打进 XPU wheel,
    # 此处以静态 grid 覆盖 (E * G) 个专家分组, token 维在核内循环。
    grid = (E * G,)
    stride_i_e, stride_i_t, stride_i_h = y.stride()
    stride_yq_e, stride_yq_t, stride_yq_h = y_q.stride()
    _silu_mul_fp8_quant_deep_gemm[grid](
        y,
        y_q,
        y_s,
        tokens_per_expert,
        H,
        group_size,
        stride_i_e,
        stride_i_t,
        stride_i_h,
        stride_yq_e,
        stride_yq_t,
        stride_yq_h,
        ys_strides[0],
        # ... 其余 stride / 形状参数 ...
    )
return y_q, y_s

```

tests/kernels/moe/test_silu_mul_fp8_quant_deep_gemm.py

测试平台化: 引入 DEVICE 常量替换硬编码 "cuda", 新增 CUDA-alike/XPU 的 skipif 门控, 并将 1-ULP FP8 容差扩展到 XPU, 使测试可在 Intel GPU 上运行并通过。

平台无关的设备常量: 所有测试张量分配统一使用当前平台设备

```
DEVICE = current_platform.device_type
```

融合内核仅对 CUDA 系 (含 ROCm) 与 XPU 有意义, 其余平台整文件跳过

```

pytestmark = pytest.mark.skipif(
    not (current_platform.is_cuda_alike() or current_platform.is_xpu()),
    reason="persistent_masked_m_silu_mul_quant requires a CUDA-alike or XPU device.",
)

```

对比断言处: ROCm 与 XPU 共用 Triton 回退内核

```
if current_platform.is_rocm() or current_platform.is_xpu():
```

Triton 的 f32 数学内建函数 (如 tl.exp) 与 torch.exp 最多差 1 ULP,

在 FP8 量化边界可能翻转一个可表示值, 允许 1 个 FP8 量子的误差。

```
torch.testing.assert_close(
    y_q[e, :nt].to(torch.float32),
    ref_y_q[e, :nt].to(torch.float32),
    atol=32.0,
    rtol=0.2,
)
else:
    torch.testing.assert_close(
        y_q[e, :nt].to(torch.float32),
        ref_y_q[e, :nt].to(torch.float32),
    )
```

评论区精华

1. oonyshch 批准时提示 please also check the path this kernel goes thru, 作者 pmanczak 回复确认: XPU 上 Triton 是唯一路径而非慢路径, 因为该 C++ 内核是 CUDA-only (activation_kernels.cu 主体在 #ifndef USE_ROCM 下), 仅在 CUDA dispatch key 下注册, 且未编译进 XPU wheel, XPU 上不存在等价 CUDA 内核可选。
 2. jikunshang 在测试中针对 device=gate.device 写法提问 why not use DEVICE?, 作者认可并统一为模块级 DEVICE 常量, 最终保持代码风格一致。
- 确认 XPU 上内核走 Triton 路径而非慢路径 (question): 确认 XPU 上无 CUDA 等价内核可选, Triton 即最优路径, 不存在性能回退风险。
 - 为何用 gate.device 而非 DEVICE 常量 (style): 作者接受建议, 测试内所有张量分配统一改用 DEVICE 常量, 保持风格一致。

风险与影响

- 风险: 低风险。主要关注点: 其一, CUDA 分支条件从 `cuda_arch >= 80` 改为 `has_device_capability(80)`, 行为等价性依赖 `has_device_capability` 内部实现与原先 `to_int()` 比较语义一致, 建议在 CUDA CI 上回归确认; 其二, 测试将 `atol=32.0 / rtol=0.2` 容差扩展到 XPU, 理论上可能掩盖真实数值偏差, 但其依据与 ROCm 相同 (Triton 回退内核的 `tl.exp` 与 `torch.exp` 最多差 1 ULP, FP8 量化边界可能翻转一个可表示值), 理由成立; 其三, `skipif` 依赖 `is_cuda_alike()` 对 ROCm 返回 True 的语义, 若未来平台分类调整需同步更新门控。
- 影响: 对 XPU (Intel GPU) 用户而言, 该融合内核路径从必然崩溃变为可正常回退 Triton 并通过测试, 运行期稳定性提升; 对 CUDA/ROCm 无运行时行为变化。对团队而言, Intel GPU CI 新增一条 FP8 MoE 量化内核路径的测试覆盖, 且测试代码平台化后更易移植到其他后端。影响范围小、程度中等偏轻。
- 风险标记: 平台判断行为依赖 `has_device_capability` 实现, CUDA 路径条件等价性需回归验证, 测试容差放宽可能掩盖数值偏差

关联脉络

- PR #51627 [Bugfix][CPU] Make the Apple Silicon BF16 probe fall back instead of raising: 同为 "平台能力探测失败时优雅回退而非断言 / 抛错" 的修复模式, 可作为同类问

题的先例对照。

- PR #50826 [XPU] [Linear] enable torch linear backend for blockwise gemm on xpu: 同为 XPU 上使能量化内核路径的工作，合起来体现 Intel GPU 后端在 vLLM 中持续补齐内核与测试覆盖。
- PR #51770 [XPU] Fix UVA weight offloading (non-pinned-tensor views and static Triton launcher): 同期 XPU 平台修复，显示 Intel 后端在持续加固，与本 PR 同属 XPU 支持演进脉络。