

PR #29027 完整报告

sgl-project/sglang

[NPU] Adding a fast layernorm for diffusion models and fix BSA

合并时间: 2026-08-05 19:06

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/29027>

执行摘要

- 一句话: NPU 扩散模型接入 fast layernorm 并修复 BSA
- 推荐动作: 建议 NPU 平台维护者精读 layernorm.py 的改造: CustomOp 注册、__init__ 内动态探测依赖、forward_* 设备拆分是清晰的适配范式, 值得作为后续 NPU kernel 接入的模板。同时注意 BSA 算子切换需与 sgl-kernel-npu 版本对齐, 并建议补充一个 forward_npu 与 forward_native 的一致性单测及 fallback 路径测试。

功能与动机

PR body 明确提到要 Support fast layernorm for diffusion models, 并关联 sgl-kernel-npu#571 的 BSA 修复。作者在 Wan2.2-T2V-A14B-Diffusers 上对比 native layernorm (146.17s) 与 fast layernorm (142.59s), 端到端加速约 1.024x。由于 NPU 上原生 F.layer_norm 无法发挥硬件能力, 需要接入 attentions 库的融合 kernel。

实现拆解

1. CustomOp 化 FP32LayerNorm (layernorm.py): 将 FP32LayerNorm 从普通 nn.LayerNorm 子类改为继承 CustomOp 并注册为 "fp32_layer_norm"。__init__ 中先调用 nn.LayerNorm.__init__ 完成参数初始化, 再通过 dispatch_forward() 按设备分派实现; 随后 try import attentions, 若不可用则打印 warning 并回退到 forward_native。这样把依赖探测放在实例化阶段, 避免顶层导入失败拖垮整个 runtime。
2. 拆分 forward 实现: 原 forward 更名为 forward_native (保留 fp32 计算语义和 _cached_fp32_param 缓存逻辑); 新增 forward_cuda 直接委托 forward_native, 保证 CUDA 行为逐位一致; 新增 forward_npu 调用 torch.ops.attentions.layernorm (impl_mode=0), 复用同一套 fp32 权重 / 偏置缓存, 保持精度语义。
3. 修复 BSA (block_sparse_attn.py): _block_sparse_attention 中把 torch.ops.attentions.block_sparse_attention 换成 torch.ops.attentions.adaptive_block_sparse_attention, 跟随 sgl-kernel-npu#571 的算子更新。
4. 配套与验证: 本 PR 未新增单元测试, 依赖 CI 集成测试 (multimodal-gen 套件、B200 与 5090 图像一致性)。B200 上 5 个 image-consistency 失败与单 GPU zimage_image_t2i_fp8 的像素指标漂移被评审判定为与 PR 无关 (PR 不影响 CUDA 路径)。

关键文件：

- python/sglang/multimodal_gen/runtime/layers/layernorm.py（模块 归一化层；类别 source；类型 core-logic；符号 FP32LayerNorm.init, FP32LayerNorm.forward_native, FP32LayerNorm.forward_cuda, FP32LayerNorm.forward_npu）：核心改动：FP32LayerNorm 从普通 nn.LayerNorm 改造为 CustomOp 注册算子，新增 forward_npu 接入 sgl-kernel-npu 的快速 layernorm，并新增 forward_cuda 委托原实现保证 CUDA 零回归；同时保留 fp32 参数缓存与动态 fallback 逻辑。
- python/sglang/multimodal_gen/runtime/layers/attention/backends/block_sparse_attn.py（模块 稀疏注意力；类别 source；类型 core-logic；符号 BlockSparseAttention._block_sparse_attention）：修复 BSA: _block_sparse_attention 中改用 ada_block_sparse_attention 算子，对应 sgl-kernel-npu#571 的更新，影响扩散模型的稀疏注意力计算路径。

关键符号：FP32LayerNorm.init, FP32LayerNorm.forward_native, FP32LayerNorm.forward_cuda, FP32LayerNorm.forward_npu, FP32LayerNorm._cached_fp32_param, BlockSparseAttention._block_sparse_attention

关键源码片段

python/sglang/multimodal_gen/runtime/layers/layernorm.py

核心改动：FP32LayerNorm 从普通 nn.LayerNorm 改造为 CustomOp 注册算子，新增 forward_npu 接入 sgl-kernel-npu 的快速 layernorm，并新增 forward_cuda 委托原实现保证 CUDA 零回归；同时保留 fp32 参数缓存与动态 fallback 逻辑。

```
# FP32LayerNorm 以 CustomOp 形式注册，sglang 会根据 inputs.device.type
# 自动分派到 forward_cuda / forward_npu；__init__ 中先探测 NPU 的
# attentions 加速库是否可用，缺失时回退到原生实现，避免启动即失败。
@CustomOp.register('fp32_layer_norm')
class FP32LayerNorm(CustomOp, nn.LayerNorm):
```

```
    def __init__(self, normalized_shape, eps=1e-5, elementwise_affine=True,
                  bias=True, device=None, dtype=None):
        # 先完成 nn.LayerNorm 初始化，保证 weight/bias 等参数照常创建
        nn.LayerNorm.__init__(
            self,
            normalized_shape=normalized_shape,
            eps=eps,
            elementwise_affine=elementwise_affine,
            bias=bias,
            device=device,
            dtype=dtype,
        )
        # dispatch_forward 依据设备选择 forward_<device> 实现
        self._forward_method = self.dispatch_forward()

    try:
        import attentions # noqa: F401 # NPU 加速库，见 sgl-kernel-npu
```

```

except ImportError:
    # 未安装 attentions 时回退原生实现, 并提示用户可安装以提速
    from sglang.multimodal_gen.runtime.utils.logging_utils import init_logger
    logger = init_logger(__name__)
    logger.warning(
        "The 'attentions' library is not installed. Falling back to "
        "native layernorm. Installing this library may improve "
        "performance on NPU. See: sgl-project/sgl-kernel-npu"
    )
    self._forward_method = self.forward_native

def _cached_fp32_param(self, attr, param, device):
    # 缓存 fp32 拷贝的参数, 保持与旧路径一致的自动求导语义;
    # 推理阶段 grad 关闭, 用 data_ptr 与 _version 作缓存键
    if param is None:
        return None
    if torch.is_grad_enabled():
        return param.float().to(device=device)
    key = (param.data_ptr(), param._version, param.device, device, param.dtype)
    cache = self.__dict__.get(attr)
    if cache is not None and cache[0] == key:
        return cache[1]
    fp32_param = param.detach().to(device=device, dtype=torch.float32)
    self.__dict__[attr] = (key, fp32_param)
    return fp32_param

def forward_native(self, inputs):
    # 原生实现: fp32 计算后转回原精度, 与 Diffusers 行为对齐
    origin_dtype = inputs.dtype
    device = inputs.device
    weight = self._cached_fp32_param('_weight_fp32_cache', self.weight, device)
    bias = self._cached_fp32_param('_bias_fp32_cache', self.bias, device)
    return F.layer_norm(
        inputs.float(),
        self.normalized_shape,
        weight,
        bias,
        self.eps,
    ).to(origin_dtype)

def forward_cuda(self, inputs):
    # CUDA 路径保持原行为, 保证本次改动不引入回归
    return self.forward_native(inputs)

def forward_npu(self, inputs):
    # NPU 路径: 调用 attentions 库的融合 kernel, impl_mode=0 为快速模式;
    # 仍复用 fp32 权重 / 偏置缓存, 保持精度语义一致
    origin_dtype = inputs.dtype
    device = inputs.device

```

```

weight = self._cached_fp32_param('_weight_fp32_cache', self.weight, device)
bias = self._cached_fp32_param('_bias_fp32_cache', self.bias, device)

output, _, _ = torch.ops.attentions.layernorm(
    input=inputs,
    normalized_shape=list(self.normalized_shape),
    weight=weight,
    bias=bias,
    eps=self.eps,
    impl_mode=0,
)
return output.to(origin_dtype)

```

python/sglang/multimodal_gen/runtime/layers/attention/backends/block_sparse_attn.py

修复 BSA: `_block_sparse_attention` 中改用 `ada_block_sparse_attention` 算子，对应 `sgl-kernel-npu#571` 的更新，影响扩散模型的稀疏注意力计算路径。

```

def _block_sparse_attention(self, query, key, value, smask, sct):
    # 切换到 ada_block_sparse_attention: 跟随 sgl-kernel-npu#571 的算子
    # 更新, 修复 BSA 在 NPU 上的行为; 参数与旧算子保持一致
    return torch.ops.attentions.ada_block_sparse_attention(
        query=query,
        key=key,
        value=value,
        sparse_mask=smask,
        sparse_count_table=sct,
        input_layout='BNSD',
        sparse_size=self.block_size,
        num_heads=query.shape[1],
        num_key_value_heads=key.shape[1],
        scale_value=self.softmax_scale,
        causal=self.causal,
        inner_precise=1,
        pre_tokens=self.default_tokens,
        next_tokens=self.default_tokens,
        actual_seq_lengths=None,
        actual_seq_lengths_kv=None,
    )

```

评论区精华

1. 顶层 import 崩溃风险 (gemini-code-assist[bot], 高优先级) : 在 `if _is_npu`: 下顶层 import `attentions` 会导致未安装该库的 NPU 平台启动即失败, 完全破坏 fallback 机制; 作者已通过 commit "Apply suggestion from @gemini-code-assist[bot]" 移除顶层导入, 改为在 `__init__` 内动态导入。
2. bare except 问题 (gemini-code-assist[bot], 中优先级) : 建议将裸 `except` 改为 `except ImportError`, 避免吞掉 `KeyboardInterrupt/SystemExit` 等异常; 已采纳。

3. fallback 提示 (ping1jing2) : 建议在回退 native 实现时打 warning, 告知用户如何安装 attentions 以及可获得的加速收益; 作者已补上 init_logger 日志。
 4. 日志工具统一 (ping1jing2) : multimodal_gen 内优先使用 `sglang.multimodal_gen.runtime.utils.logging_utils.init_logger`, 作者回复 done。
- 顶层 import attentions 破坏 fallback 机制 (correctness): 作者通过 commit "Apply suggestion from @gemini-code-assist[bot]" 移除顶层导入, 动态导入保留在 `__init__` 中。
 - bare except 应改为 except ImportError (style): 已修改为 `except ImportError`。
 - fallback 时补充 warning 日志引导安装 (documentation): 作者已补充 `init_logger` warning, 提示安装 `sgl-kernel-npu` 可提升 NPU 性能。
 - 日志工具统一使用 `logging_utils.init_logger` (style): 作者回复 done, 已切换日志工具。

风险与影响

- 风险:
 1. 外部依赖运行时错误: `forward_npu` 强依赖 `sgl-kernel-npu` 的 `attentions` 包, fallback 只在 `ImportError` 时生效; 若包已安装但版本不匹配或算子签名变化, 会在运行时报错且无降级。
 2. BSA 算子跨平台影响: `block_sparse_attn.py` 的改动影响所有使用该 backend 的设备, 若 CUDA 侧没有 `ada_block_sparse_attention` 或语义不一致, 可能引入回归; 需要与 `sgl-kernel-npu#571` 版本严格配套。
 3. 缺少单元测试: `forward_npu` 与 `forward_native` 的一致性、fallback 路径均无单测覆盖; CI 中 B200 图像一致性失败虽被判定无关, 但像素级指标漂移 ($SSIM\ 0.8471 < 0.95$) 仍需持续关注。
 4. 性能收益有限: 端到端加速仅 1.024x, 若用户在非瓶颈场景使用, 收益感知不强, 但 layernorm 单算子收益可能更高。- 影响: 用户侧: NPU (昇腾) 用户安装 `sgl-kernel-npu` 后, 扩散模型 (Wan2.2、FLUX 等) 的 `FP32LayerNorm` 自动走快速 kernel; 未安装则回退原生实现并收到 warning 引导。系统侧: 改动集中在 `multimodal_gen runtime` 的 layernorm 与 BSA backend, CUDA 等平台行为完全不变, 改动量小 (+57/-3), 风险面可控。团队侧: 确立了 CustomOp 注册 + 设备分派 + 动态导入 fallback 的 NPU 接入模板, 后续 NPU kernel 可复用该模式。- 风险标记: 外部依赖 `sgl-kernel-npu`, 缺少单测覆盖, BSA 算子跨平台切换, 性能收益有限

关联脉络

- PR #30883 [XPU] Add qknorm_rope support for Flux: 同样修改 `layernorm.py` 并为特定平台 (XPU) 定制 diffusion 算子, 与本 PR 构成逐平台适配的平行演进线。
- PR #33599 [AMD] Fuse Kimi-K3 attn-residual aggregation: 平台特定 kernel 融合与接入范式相似, 体现了 `sglang` 多后端 kernel 适配的统一思路。
- PR #28040 [Intel GPU] DeepSeek V4 8/N: use sgl-kernel implementation of fused_k_norm_rope_flashmla on XPU: 同为将平台加速 kernel 接入 `sglang` 层的做法, 与本 PR 的 NPU kernel 接入模式一致。

- PR #33523 [npu] [bugfix] Fix PD-disaggregation error: 同一 NPU 平台维护线的 bugfix , 说明 NPU 适配正在持续完善。