

PR #37162 完整报告

sgl-project/sglang

[Diffusion] Fuse FLUX.2 ModelOpt FP8 producers and QKV packing

合并时间: 2026-09-01 16:14

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/37162>

执行摘要

- 一句话: FLUX.2 FP8 模型中融合 LayerNorm、量化、QKV 打包等核心算子, 显著降低内核数量并提升推理速度。
- 推荐动作: 此 PR 非常值得精读。它是一个高质量、低风险的性能优化范例: 动机明确, 实现详尽且有充分数据支撑, 测试覆盖全面, 并设计了安全的回退机制。值得关注的设计决策包括: 1) 基于 BitExactFusionGate 的位精确验证与自动回退模式; 2) 精细的作用域控制 (硬件、TP、编译器状态); 3) 将性能分析从宏观 (e2e) 到微观 (内核计数、微基准测试) 的完整呈现。对于从事深度学习推理优化、内核开发或 GPU 编程的工程师, 该 PR 的实现模式和技术细节具有很高的参考价值。

功能与动机

实现 FLUX.2 ModelOpt FP8 路径上剩余的性能优化, 参考《Agentic Kernels in Production》文章中的方法。目标是减少内核启动开销、中间内存分配和数据搬运, 从而提升整体推理速度。PR body 中明确指出, 完整路径达到 -3.20% e2e / -3.25% denoise 的性能提升。

实现拆解

1. 添加 FP8 融合门控与线性层验证逻辑: 在 `python/sglang/multimodal_gen/runtime/models/dits/flux_2.py` 中, 新增 `_valid_modelopt_fp8_linear` 和 `_shared_modelopt_fp8_scale` 函数, 用于验证 FP8 线性层的 `input_scale` 属性是否满足融合条件。同时引入新的 `BitExactFusionGate (_FLUX2_LN_FP8)` 以管理 LN+modulate+FP8 量化的融合路径。
2. 创建融合的 LayerNorm 调制与 FP8 量化内核: 在 `python/sglang/kernels/ops/diffusion/norm/layernorm_modulate_triton.py` 中, 新增 `fused_layernorm_modulate_fp8_quant_raw` 函数。该内核将原有的 `fused_layernorm_modulate_raw` 与 `static_quant_fp8` 操作合并, 直接在 GPU 上完成归一化、仿射变换和 FP8 量化, 消除了中间 BF16 张量的读写。在模型主文件中, 通过 `_try_flux2_norm_modulate_fp8` 函数作为入口, 在满足条件时调用此融合内核。
3. 合并 QKV GEMM 与添加 QKV epilogue JIT 内核: 修改 `Flux2TransformerBlock` 的初始化逻辑, 当检测到 ModelOpt FP8 配置、TP=1 且在 Blackwell (sm >= 10) 架构上时, 启用 `use_fused_qkv` 和 `use_fused_qkv_epilogue`。这会将 image/text 分离的 Q、K、V 投影合并为两个 channelwise-CUTLASS GEMM。随后在 forward 路径中, 通过新增的 `try_fused_flux2_qkv_epilogue` 函数 (位于 `python/sglang/kernels/ops/diffusion/rope/flux2_qkv_epilogue_jit.py`), 将 QK RMSNorm、RoPE、QKV 打包和文本 / 图像拼接融合为

一个 Blackwell JIT 内核。

4. 融合 token 拼接与 FP8 量化：在 `python/sglang/kernels/ops/diffusion/layout/flux2_token_cat_fp8_triton.py` 中，新增 `try_flux2_token_cat_fp8` 函数和对应的 Triton 内核 `_token_cat_fp8_kernel`。该融合操作将 single-block 中 attention 与 MLP 的拼接结果直接静态量化为 FP8，避免生成全宽度的 BF16 中间张量。
5. 广泛的测试与基准测试配套：新增了多个单元测试文件（如 `test_flux2_fp8_norm_quant_gate.py`, `test_flux2_qkv_epilogue.py`, `test_flux2_token_cat_fp8.py`, `test_flux2_fp8_norm_quant.py`）来验证各个融合内核的位精确性、对编译器和 CUDA Graph 的正确降级行为。同时新增了基准测试文件（如 `bench_flux2_fp8_norm_quant.py`, `bench_flux2_token_cat_fp8.py`）。修改了现有测试（如 `test_transformer_quant.py`, `test_component_accuracy_weight_transfer.py`）以覆盖新的合并 QKV 权重格式和条件权重迁移。所有变更均通过现有测试套件（2372 passed）。

关键文件：

- `python/sglang/multimodal_gen/runtime/models/dits/flux_2.py`（模块 扩散模型；类别 source；类型 core-logic；符号 `_valid_modelopt_fp8_linear`, `_shared_modelopt_fp8_scale`, `_try_flux2_norm_modulate_fp8`, `_flux2_norm_maybe_fp8`）：核心模型文件，实现了 FP8 融合路径的接入逻辑、条件配置和 forward 中的调度，是本次优化的主体。
- `python/sglang/kernels/ops/diffusion/rope/flux2_qkv_epilogue_jit.py`（模块 内核库；类别 infra；类型 infrastructure；符号 `try_fused_flux2_qkv_epilogue`, `flux2_qkv_epilogue_module`）：新增的 QKV epilogue JIT 内核实现，是本次性能提升的关键内核之一。
- `python/sglang/kernels/ops/diffusion/layout/flux2_token_cat_fp8_triton.py`（模块 内核库；类别 infra；类型 infrastructure；符号 `try_flux2_token_cat_fp8`, `_token_cat_fp8_kernel`）：新增的 token 拼接与 FP8 量化的融合 Triton 内核，避免了生成全宽度 BF16 中间张量。

关键符号：`_valid_modelopt_fp8_linear`, `_shared_modelopt_fp8_scale`, `_try_flux2_norm_modulate_fp8`, `_flux2_norm_maybe_fp8`, `configure_fp8_norm_quant`, `try_fused_flux2_qkv_epilogue`, `try_flux2_token_cat_fp8`, `fused_layernorm_modulate_fp8_quant_raw`

关键源码片段

`python/sglang/multimodal_gen/runtime/models/dits/flux_2.py`

核心模型文件，实现了 FP8 融合路径的接入逻辑、条件配置和 forward 中的调度，是本次优化的主体。

以下为 `flux_2.py` 中新增的关键辅助函数与配置方法的核心逻辑

验证单个 ModelOpt FP8 线性层的 `input_scale` 属性是否满足融合条件

```
def _valid_modelopt_fp8_linear(linear: nn.Module) -> bool:
    input_scale = getattr(linear, "input_scale", None)
    return (
```

```

        isinstance(getattr(linear, "quant_method", None), ModelOptFp8LinearMethod)
        and isinstance(input_scale, torch.Tensor)
        and input_scale.is_cuda
        and input_scale.dtype == torch.float32
        and input_scale.numel() == 1
        and input_scale.is_contiguous()
        and bool(torch.isfinite(input_scale).all().item())
        and bool((input_scale > 0).all().item())
    )

# 尝试融合 LayerNorm + adaLN 调制 + 静态 FP8 量化, 返回融合后的张量或 None
# 此函数包含了严格的形状、签名检查和位精确验证逻辑
try:
    out = fused_layernorm_modulate_fp8_quant_raw(
        x, scale_row, shift_row, input_scale, norm.eps
    )
except Exception as exc:
    _FLUX2_LN_FP8.on_exception(exc, logger=logger)
    return None
# ... ( 与参考实现进行位精确比较 ) ...
return _FLUX2_LN_FP8.accept_or_fallback(
    out,
    reference,
    sig=sig,
    logger=logger,
    mismatch_msg="...",
)

# 在 TransformerBlock 初始化中, 配置合并 QKV 与 epilogue
fp4_packed_qkv = isinstance(quant_config, ModelOptFp4Config) and getattr(
    quant_config, "checkpoint_uses_packed_qkv", False
)

capability = current_platform.get_device_capability()
# 仅当 ModelOpt FP8 配置、TP=1 且为 Blackwell (sm >= 10) 架构时启用合并 QKV
fp8_merged_qkv = (
    isinstance(quant_config, ModelOptFp8Config)
    and self.tp_size == 1
    and capability is not None
    and capability.major >= 10
)

self.use_fused_qkv = fp4_packed_qkv or fp8_merged_qkv
self.use_fused_qkv_epilogue = fp8_merged_qkv

```

python/sglang/kernels/ops/diffusion/rope/flux2_qkv_epilogue_jit.py

新增的 QKV epilogue JIT 内核实现, 是本次性能提升的关键内核之一。

```

# 尝试执行融合的 QK 归一化、RoPE、QKV 打包和文本 / 图像拼接内核
def try_fused_flux2_qkv_epilogue(
    img_q: torch.Tensor, # 图像 Q 张量 ( 非连续, 来自合并 GEMM 的视图 )

```

```

img_k: torch.Tensor,
img_v: torch.Tensor,
txt_q: torch.Tensor, # 文本 Q 张量
txt_k: torch.Tensor,
txt_v: torch.Tensor,
img_q_weight: torch.Tensor, # 图像 Q/RMSNorm 权重
img_k_weight: torch.Tensor,
txt_q_weight: torch.Tensor, # 文本 Q/RMSNorm 权重
txt_k_weight: torch.Tensor,
cos_sin_cache: torch.Tensor, # 预计算的 RoPE cos/sin 缓存
img_eps: float, # 归一化 epsilon
txt_eps: float,
) -> tuple[torch.Tensor, torch.Tensor, torch.Tensor] | None:
    # 严格的设备、形状、对齐和编译器状态检查, 不满足则返回 None 回退
    if torch.compiler.is_compiling():
        return None
    if not ( ... ): # 检查张量属性、Blackwell 架构、非 CUDA Graph 捕获等
        return None
    # ... ( 验证权重和缓存属性 ) ...
    # 计算总 token 数并分配输出缓冲区
    total_tokens = txt_q.shape[1] + img_q.shape[1]
    joint_shape = (1, total_tokens, heads, _HEAD_DIM)
    joint_q = torch.empty(joint_shape, dtype=img_q.dtype, device=img_q.device)
    # ... (joint_k, joint_v 类似) ...
    # 调用 JIT 编译的 CUDA 内核执行融合操作
    flux2_qkv_epilogue_module().flux2_qkv_epilogue(
        joint_q.view(-1, heads, _HEAD_DIM),
        joint_k.view(-1, heads, _HEAD_DIM),
        joint_v.view(-1, heads, _HEAD_DIM),
        img_q.view(-1, heads, _HEAD_DIM),
        # ... ( 其他参数 ) ...
    )
    return joint_q, joint_k, joint_v

```

python/sglang/kernels/ops/diffusion/layout/flux2_token_cat_fp8_triton.py

新增的 token 拼接与 FP8 量化的融合 Triton 内核，避免了生成全宽度 BF16 中间张量。

后端 Triton 内核：在同一个 kernel 中完成拼接和 FP8 量化

@triton.jit

```

def _token_cat_fp8_kernel(
    attention, # 输入 attention 张量
    mlp, # 输入 mlp 张量
    output, # 输出 FP8 张量 (uint8)
    input_scale, # 静态量化 scale
    attention_hidden: tl.constexpr, # attention 的 hidden dim
    mlp_hidden: tl.constexpr, # mlp 的 hidden dim
    output_hidden: tl.constexpr, # 输出的 hidden dim
    BLOCK: tl.constexpr, # Triton block size
    FP8_DTYPE: tl.constexpr, # FP8 数据类型

```

```

FP8_MIN: tl.constexpr, # FP8 最小值
FP8_MAX: tl.constexpr, # FP8 最大值
USE_PDL: tl.constexpr, # 是否使用 PDL 加速
):
    row = tl.program_id(0)
    block = tl.program_id(1)
    columns = block * BLOCK
    # ... ( 计算列索引和掩码 ) ...
    # 根据列索引决定从 attention 还是 mlp 加载数据
    attention_values = tl.load(attention + row * attention_hidden + columns, mask=attention_
    mask, ...).to(tl.float32)
    mlp_values = tl.load(mlp + row * mlp_hidden + mlp_columns, mask=mlp_mask, ...).to(tl.
    float32)
    values = tl.where(columns < attention_hidden, attention_values, mlp_values)
    # 执行静态 FP8 量化
    scale = tl.load(input_scale).to(tl.float32)
    quantized = tl.clamp(values * (1.0 / scale), FP8_MIN, FP8_MAX).to(FP8_DTYPE)
    # 写入输出
    tl.store(output + row * output_hidden + columns, quantized.to(tl.uint8, bitcast=True), mask=
    output_mask)

# 前端入口函数, 包含完整的参数校验和内核启动逻辑
def try_flux2_token_cat_fp8(attention: torch.Tensor, mlp: torch.Tensor, input_scale: torch.
Tensor) -> torch.Tensor | None:
    # ... ( 检查张量属性、Blackwell 架构、非编译器 /Graph 状态 ) ...
    output = torch.empty((*attention.shape[:-1], output_hidden), dtype=fp8_dtype, device=
    attention.device)
    # 启动 Triton 内核
    _token_cat_fp8_kernel[(rows, triton.cdiv(output_hidden, _BLOCK))](
        attention, mlp, output.view(torch.uint8), input_scale,
        # ... ( 其他常量参数 ) ...
    )
    return output

```

评论区精华

PR body 中详细阐述了设计决策与权衡：

- 作用域与降级路径：明确指出所有优化仅在 Blackwell, TP=1 上启用。对于 BF16、Hopper FP8、TP>1、torch.compile、CUDA Graph 捕获、不支持的布局以及文本序列尾部填充等情况，均回退到现有的分离路径。Packed NVFP4 行为不受新 QKV epilogue 影响。
- 性能与准确性的平衡：提供了详尽的端到端基准测试（GB300 硬件，1024x1024，50 步），证明 PR 分支与 main 分支生成像素完全一致（PSNR 无穷大，SSIM 1.0），性能提升 -3.20%。同时展示了生产者内核的微基准测试，token-cat+FP8 量化内核获得 1.77x-2.60x 加速。
- 内核数量分析：通过 Torch profiler 数据证明，总内核数从 7335 降至 5878 (-19.9%)，其中静态 FP8 量化内核减少 83.4%，cat 内核减少 53.8%。端到端 GPU 时间的提升主要来自消除了这些周围的生产者、拷贝和启动开销。

- 精确验证与回退机制：所有新增的融合路径均通过 `BitExactFusionGate` 的签名验证机制和运行时与参考实现的位精确比较进行守护，确保在任何平台或形状下，若融合结果不位精确则自动回退，保证了正确性。
- 暂无高价值评论线程

风险与影响

- 风险：
 1. 正确性风险：融合内核（特别是 JIT 内核）必须在所有支持的硬件和输入形状下保持与分离路径的位精确一致。PR 通过严格的门控签名验证、运行时回退机制以及广泛的单元测试（位精确测试、边界形状测试）和端到端图像对比测试来缓解此风险。测试覆盖了 1、127、4096 token 等多种情况。
 2. 性能风险：融合操作（如更宽的合并 GEMM）自身并非在所有情况下都更快。PR body 指出，宽 GEMM 的聚合时间（501ms -> 510ms）略有增加，但整体收益来自消除周围操作。主要风险在于，在非目标平台或异常输入下，可能意外触发融合路径导致性能下降，但现有回退逻辑应能处理。
 3. 兼容性风险：引入了新的合并 QKV 权重格式。需要确保模型加载器能正确处理 FP8 格式下合并与分离的 QKV 投影，并且现有权重转换工具（`AccuracyEngine.transfer_weights`）能适配条件映射。PR 中已通过修改 `test_component_accuracy_weight_transfer.py` 测试覆盖了此场景。
 4. 可维护性风险：新增了多个 JIT 内核（CUDA 和 Triton）和复杂的门控逻辑，增加了代码库的复杂性。需要团队熟悉这些新的融合模式和验证机制。
 - 影响：用户影响：在 Blackwell (GB300) GPU 上运行 FLUX.2 FP8 模型的用户将获得约 3.2% 的端到端推理速度提升和约 438 MB 的显存节省。生成质量（像素级）与优化前完全一致。对于使用其他硬件或配置的用户，行为无变化（自动降级）。系统影响：显著减少了推理过程中的 CUDA 内核调用数量（约 20%），降低了内核启动开销和 CPU-GPU 同步压力。引入了新的 JIT 编译内核，增加了对 Blackwell 架构特定优化的依赖。可能与其他扩散模型（如 FLUX.1）或未来的 FP8 模型产生协同优化模式。团队影响：需要扩散模型和内核开发团队维护这些新增的融合内核和复杂的门控逻辑。详细的基准测试和分析方法为后续性能优化提供了范例。测试套件的扩充有助于保证质量。
- 风险标记：核心路径变更，平台特定优化，需要严格测试覆盖

关联脉络

- PR #37112 [Diffusion] Fuse FLUX.2 gated residual normalization on Blackwell: 同为 FLUX.2 模型的 Blackwell JIT 内核融合优化，是本 PR 优化的前序工作（融合了残差归一化）。
- PR #37123 [Diffusion] Fuse Qwen-Image FP8 QKV projection and Blackwell epilogue: 同属 Diffusion 子系统 FP8 融合优化系列，为不同模型（Qwen-Image）实现了类似的 FP8 QKV 投影与 epilogue 融合模式。
- PR #37129 [Diffusion] Fuse Qwen-Image residual norm and NVFP4 quantization: 同属 Diffusion 子系统内核融合优化，展示了将归一化与量化操作融合的相似技术路径。

- PR #37299 refactor(hicache): simplify decode offload state bookkeeping: 同为性能与正确性优化，虽属不同子系统（hicache vs diffusion），但都体现了通过简化状态管理和增加测试来提升系统健壮性的思路。