

PR #33702 完整报告

sgl-project/sglang

[diffusion] Add Sol-Attn sparse attention backend for diffusion

合并时间: 2026-08-09 16:26

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/33702>

执行摘要

- 一句话: 新增 Sol-Attn 稀疏注意力后端, H3 去噪提速 1.15-1.23x
- 推荐动作: 值得精读 `sol_attn.py` 与 `cuda.py`: 展示了可选三方内核接入 `sglang` 的标准模式 (枚举注册、懒加载 `resolver`、`dense guard` 兜底), 以及 `diffusion` 场景下 "质量 - 性能" 权衡的参数化设计。对想扩展 `diffusion` 稀疏注意力后端或理解注意力后端插件机制的开发者有直接参考价值。建议重点关注 `_should_use_dense` 的层 / 步混合门控与 `_resolve_kv_splits` 的硬件自适应逻辑。

功能与动机

MiniMax-H3 的 denoise 是 attention-heavy 负载: 约 50 层、每步约 38k packed tokens 全部经由 `forward_varlen` 走 FlashAttention (FA3)。Sol-Attn 是 NVLabs 的 training-free 稀疏注意力方法, 在 `online-softmax` 过程中按阈值路由 KV 块, 无需物化完整代理分数图即可跳过大量计算。PR 目标即把它注册为标准 `diffusion` 注意力后端, 默认配置下用约 15% 的 denoise 提速换取 32.9 dB 的 PSNR (t2va 场景), 并沉淀出可供后续模型复用的后端扩展样板。

实现拆解

1. 注册枚举与稀疏标记: 在 `python/sglang/multimodal_gen/runtime/platforms/interface.py` 的 `AttentionBackendEnum` 中新增 `SOL_ATTEN`, 并加入 `is_sparse` 属性集合, 使平台层将该后端识别为稀疏注意力家族成员。
2. 新增核心后端: 新建 `python/sglang/multimodal_gen/runtime/layers/attention/backends/sol_attn.py`, 实现 `SolAttnBackend` (约束 `head_size=128`, 暴露 `get_enum()`、`get_impl_cls()`) 与 `SolAttnImpl`。关键设计是 `_should_use_dense()` 的混合门控: 前 `dense_steps` 个去噪步与 `dense_layers` 指定的层保持稠密 FA 计算, 其余走 `_run_sol_attn_thd` 调上游 `sol_attn` 内核; `_resolve_kv_splits` 在 H200 (计算能力 9.0) 且序列不低于 65536 token 时自动把 `kv_splits` 解析为 4, 其余情况回退为 1。
3. 平台接线: 在 `python/sglang/multimodal_gen/runtime/platforms/cuda.py` 新增 `_SolAttnBackendResolver`, 采用懒加载 `import` 模式: 三方包未安装时报出安装指引并抛 `ImportError`, 避免在 `import` 期崩溃; 随后注册进 `_CUDA_ATTENTION_BACKEND_RESOLVERS` 映射。
4. H3 接入: 在 `python/sglang/multimodal_gen/runtime/models/dits/minimax_h3.py` 的 `MiniMaxH3Attention` 中保存 `self.prefix`, 并在 `_set_attention_backend` 构造

SolAttnImpl 时透传 prefix, 使后端能用正则 blocks.(\d+) 解析层号并正确执行 dense_layers 匹配。

5. 测试与文档配套: 新增 python/sglang/multimodal_gen/test/unit/test_sol_attn_backend.py, 覆盖枚举名、_parse_layer_ranges、head size 约束、dense guard 的 step 与 layer 判定以及 CUDA resolver 解析; 在 docs/docs/sglang-diffusion/attention_backends.mdx 补齐安装命令、参数表与平台支持矩阵。

配套说明: 上游 sol-attn 包不随 sglang 分发, 需手动 pip install

git+https://github.com/NVlabs/Sana.git@sol-engine#subdirectory=techniques/sparse_backends; CI 未安装该包时 test_cuda_resolver 自动 skipTest, 实际回归主要依赖手动基准。

开发过程中曾接入 Wan 自注意力, 但因 PSNR 仅 8.2 dB 被移除。

关键文件:

- python/sglang/multimodal_gen/runtime/layers/attention/backends/sol_attn.py (模块 注意力后端; 类别 source; 类型 core-logic; 符号 SolAttnBackend, SolAttnImpl, _parse_layer_ranges, _resolve_kv_splits) : PR 核心: 新增 SolAttnBackend 与 SolAttnImpl, 实现 packed varlen 稀疏注意力路径、dense guard 混合门控与运行时配置解析。
- python/sglang/multimodal_gen/test/unit/test_sol_attn_backend.py (模块 单元测试; 类别 test; 类型 test-coverage; 符号 FakeCudaPlatform, TestSolAttnBackend, test_parse_layer_ranges, test_dense_guard_uses_early_steps) : 新增单元测试, 覆盖 dense guard 的 step/layer 判定、layer range 解析与 resolver 解析, 是验证可选后端行为的关键配套。
- python/sglang/multimodal_gen/runtime/platforms/cuda.py (模块 平台路由; 类别 source; 类型 dependency-wiring; 符号 _SolAttnBackendResolver, resolve) : 平台接线关键改动: 新增 _SolAttnBackendResolver 并将 SOL_ATTEN 注册进 CUDA 注意力后端解析表, 决定 CLI 参数如何映射到实现类。
- python/sglang/multimodal_gen/runtime/models/dits/minimax_h3.py (模块 模型接入; 类别 source; 类型 data-contract; 符号 MiniMaxH3Attention, _set_attention_backend) : H3 接入的唯一改动: 保存 prefix 并透传给注意力实现, 使 Sol-Attn 的 dense_layers 层号匹配成为可能, 是 review 讨论的焦点。
- python/sglang/multimodal_gen/runtime/platforms/interface.py (模块 接口枚举; 类别 source; 类型 core-logic; 符号 AttentionBackendEnum, SOL_ATTEN, is_sparse) : 枚举层登记: 新增 AttentionBackendEnum.SOL_ATTEN 并标记为 is_sparse, 是后端可被 CLI 与平台调度识别的入口。
- docs/docs/sglang-diffusion/attention_backends.mdx (模块 使用文档; 类别 docs; 类型 documentation) : 用户文档: 补充 sol_attn 的 CLI 名、安装命令、全部配置参数表 (tau/thresh_type/sink_tokens/dense_steps 等) 与平台支持矩阵。

关键符号: SolAttnBackend.get_supported_head_sizes, SolAttnBackend.get_enum, SolAttnBackend.get_impl_cls, SolAttnImpl._should_use_dense, SolAttnImpl._dense_varlen, SolAttnImpl._run_sol_attn_thd, SolAttnImpl.forward, SolAttnImpl.forward_varlen, _parse_layer_ranges, _resolve_kv_splits, _get_sol_attn_runtime_config, _SolAttnBackendResolver.resolve

关键源码片段

python/sglang/multimodal_gen/runtime/platforms/cuda.py

平台接线关键改动：新增 `_SolAttnBackendResolver` 并将 `SOL_ATTN` 注册进 CUDA 注意力后端解析表，决定 CLI 参数如何映射到实现类。

```
class _SolAttnBackendResolver(_CudaAttentionBackendResolver):
    backend = AttentionBackendEnum.SOL_ATTN

    @classmethod
    def resolve(cls, platform) -> str:
        try:
            # 懒加载：sol_attn 是可选三方依赖，未安装时只报错并给出指引，
            # 避免在 import 期直接崩溃，保证 CI 与未安装用户不受影响
            from sol_attn import sol_attn # noqa: F401
            from sglang.multimodal_gen.runtime.layers.attention.backends.sol_attn import ( # noqa: F401
                SolAttnBackend,
            )

            return "sglang.multimodal_gen.runtime.layers.attention.backends.sol_attn.
                SolAttnBackend"
        except ImportError as e:
            logger.error("Failed to import Sol-Attn backend: %s", str(e))
            raise ImportError(
                "Sol-Attn backend is not installed. Install it with "
                "`pip install git+https://github.com/NVlabs/Sana.git@sol-engine"
                "#subdirectory=techniques/sparse_backends`."
            ) from e
```

评论区精华

review 中唯一实质性线程围绕 `prefix` 透传的作用域展开：

- mickqian (reviewer) 在 `minimax_h3.py:580` 提问："does other model needs modifications as well?", 质疑是否只改 H3 就够了。
- niehen6174 (作者) 回应："H3-specific change. H3 uses a custom lazy attention path (forward_varlen) that didn't pass prefix before; Sol-Attn needs it for dense_layers. Models using USPAttention already pass prefix, so no change needed there. Other models' `_supported_attention_backends` whitelist has not been updated yet ... That will be done in a follow-up PR."

结论：线程已解决——作者确认该改动为 H3 专属，其他模型的白名单扩展明确推迟到 follow-up PR；mickqian 随后以 APPROVED 收尾。

- `prefix` 透传是否为 H3 专属改动，其他模型是否需要同步修改 (question): 确认 H3 专属改动；多模型白名单扩展推迟到 follow-up PR，mickqian 随后 APPROVED。

风险与影响

- 风险：
 - 依赖可复现性：sol_attn 内核来自 NVlabs/Sana 的 sol-engine 分支且未固定 commit，安装命令可能随时间漂移；CI 在该依赖缺失时自动跳过测试，回归只能依赖手动基准与 nightly。
 - 精度敏感性：tau、dense_steps 对质量影响陡峭（PR 自测从默认 32.9 dB 一路掉到 19.3 dB），后端未对配置范围做校验，误用极易产生低质量输出；建议只使用文档默认配置。
 - 范围假设硬失败：SolAttnImpl 对非 128 head size 抛 ValueError、对非 bf16 输入抛 TypeError，属于安全但严格的上限约束；_resolve_kv_splits 依赖 torch.cuda.get_device_capability，无 CUDA 环境不可用（后端本身声明为 CUDA-only）。
 - 回归面：minimax_h3.py 的 _set_attention_backend 改动对所有注意力后端生效，但 prefix 参数带默认值，兼容性风险低；默认路径仍为 FA3，未显式选择 sol_attn 的用户完全不受影响。
 - CI 状态：合并时 pr-test-extra 运行标记为失败，虽与可选依赖缺失有关，但合并前未完全清零，值得在后续 PR 中补充跳过逻辑的显式验证。
- 影响：
 - 用户侧：H3 用户在 H200 上可显式开启后端获得 1.15-1.23x denoise 加速、PSNR 31-33 dB；默认体验不受影响。
 - 系统侧：diffusion 平台注册表新增一个三方稀疏注意力后端，cuda.py 的 resolver 表、文档参数矩阵与 CI 测试注册表需要同步维护。
 - 团队侧：为后续稀疏后端接入（MoBA、SLA 等）提供了标准样板，同时积累了对 Sol-Attn 质量 - 性能权衡的实测数据；Wan 的 Morton 重排支持、eager CuTe warmup、多模型白名单扩展已列入 route map。
 - 风险标记：可选第三方依赖未随包分发，精度敏感参数无范围校验，仅 H3 单模型验证，CI 未覆盖 sol_attn 包，pr-test-extra 合并时为失败状态

关联脉络

- PR #34107 [diffusion] fix: guard SageAttention SM90 bindings: 同样修改 cuda.py 的注意力后端 resolver 区域，说明 diffusion 后端注册表正在快速扩展，多个三方内核以相同懒加载模式接入。
- PR #34085 [diffusion] Clean up kernels and shared fast paths: 集中 diffusion 共享快路径与质量门控，与本 PR 的 dense guard 兜底设计同属同一 "快路径 + 质量兜底" 体系。
- PR #34126 [diffusion] FLUX.1: route the adaLN LN+modulate sites through the bit-exact fused LayerNorm+modulate kernel: 同属 diffusion 去噪性能优化系列，体现 diffusion 后端性能演进方向，本 PR 是其中稀疏注意力方向的最新一环。