

PR #27469 完整报告

sgl-project/sglang

dflash add sliding window attention draft layer support

合并时间: 2026-06-14 15:32

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/27469>

执行摘要

- 一句话: 为 DFlash 草稿模型添加滑动窗口注意力层支持
- 推荐动作: 值得精读, 尤其是 flashinfer_backend.py 中处理 prefix_lens 为 None 的逻辑设计, 以及 DFlash verify 路径直接裁剪 kv_start_idx 的权衡。测试覆盖了基本场景, 但建议后续补充非 FlashInfer 后端验证。

功能与动机

PR 从 #23000 拆分, 旨在为 DFlash 草稿模型提供滑动窗口注意力层支持, 使草稿模型可以按层选择高效滑动窗口注意力, 减少 KV 缓存占用并提升性能。

实现拆解

1. 工具函数 (dflash_utils.py): 新增 get_dflash_layer_types 和 get_dflash_attention_sliding_window_size, 从模型配置中读取 layer_types 序列和滑动窗口大小 (HF 风格转 SGLang window_left)。
2. 模型层 (dflash.py): 新增 _get_dflash_layer_attention_params, 根据层类型返回 (sliding_window_size, AttentionType); 在 DFlashAttention 中调用并将结果传入 RadixAttention, 替代硬编码的 ENCODER_ONLY。同时新增 get_attention_sliding_window_size 方法供外部使用。
3. 注意力后端 (flashinfer_backend.py): 修改 update_sliding_window, 当 prefix_lens 为 None (来自 target-verify) 时, 从 spec_info 的 num_accept_tokens 推算 prefix_lens; 修改 call_begin_forward, 对 DFLASH_VERIFY 类型直接使用 kv_start_idx 裁剪 KV 索引, 避免自定义掩码。
4. 测试框架 (speculative_target_verify_runner.py): 新增 _make_flashinfer_dflash_swa_built_in_masks 和 _expected_case_and_masks_for_spec_verify, 模拟 FlashInfer 内置掩码生成; 调整 _make_spec_verify_input 在满足条件时置空 custom_mask。
5. 测试用例 (test_flashinfer.py): 新增 SPEC_VERIFY_CASES 和 SPEC_VERIFY_CUDA_GRAPH_CASES, 包含 sliding_window_size=4 的 DFLASH 验证链场景。
6. 文档 (README.md): 更新覆盖矩阵, 将 FlashInfer SWA 验证路径从 blocked 改为支持 DFLASH TARGET_VERIFY。

关键文件：

- python/sglang/srt/models/dflash.py（模块 模型定义；类别 source；类型 data-contract；符号 `_get_dflash_layer_attention_params`, `get_attention_sliding_window_size`）：核心模型文件中新增 `_get_dflash_layer_attention_params` 和 `get_attention_sliding_window_size`，修改 DFlashAttention 初始化为动态注意力类型，是整个变更的入口。
- python/sglang/srt/speculative/dflash_utils.py（模块 工具函数；类别 source；类型 dependency-wiring；符号 `get_dflash_layer_types`, `get_dflash_attention_sliding_window_size`）：新增两个关键工具函数，为 dflash.py 和外部查询提供配置读取能力。
- python/sglang/srt/layers/attention/flashinfer_backend.py（模块 注意力后端；类别 source；类型 dependency-wiring）：修改滑动窗口更新器和 prefill 元数据生成，使 target-verify 路径能支持 SWA；是注意力后端的核心适配。
- python/sglang/test/kits/attention_unittest/runner_modes/speculative_target_verify_runner.py（模块 测试框架；类别 test；类型 test-coverage；符号 `_make_flashinfer_dflash_swa_built_in_masks`, `_expected_case_and_masks_for_spec_verify`）：测试框架核心文件，新增 FlashInfer 内置掩码模拟函数和预期调整函数，是测试可重复性的基础。
- test/registered/attention/unittests/swa/test_flashinfer.py（模块 单元测试；类别 test；类型 test-coverage；符号 `test_runner_mode_spec_verify_cases`, `test_runner_mode_spec_verify_cuda_graph_cases`）：新增 SWA 验证链的测试用例和 CUDA 图测试，确保新功能在 FlashInfer 后端正确运行。
- test/registered/attention/unittests/swa/README.md（模块 文档；类别 docs；类型 documentation）：更新覆盖矩阵，说明 FlashInfer SWA 现在支持 DFLASH TARGET_VERIFY，反映功能状态。

关键符号：`_get_dflash_layer_attention_params`, `get_dflash_layer_types`, `get_dflash_attention_sliding_window_size`, `FlashInferIndicesUpdaterPrefill.update_sliding_window`, `FlashInferIndicesUpdaterPrefill.call_begin_forward`, `_make_flashinfer_dflash_swa_built_in_masks`, `_expected_case_and_masks_for_spec_verify`, `test_runner_mode_spec_verify_cases`, `test_runner_mode_spec_verify_cuda_graph_cases`

关键源码片段

python/sglang/srt/models/dflash.py

核心模型文件中新增 `_get_dflash_layer_attention_params` 和 `get_attention_sliding_window_size`，修改 DFlashAttention 初始化为动态注意力类型，是整个变更的入口。

```
# python/sglang/srt/models/dflash.py
```

```
# 新增函数：根据配置和层 ID 返回滑动窗口大小和注意力类型
def _get_dflash_layer_attention_params(
    config, layer_id: int
```

```

) -> Tuple[int, AttentionType]:
    # 从配置中读取所有层的类型序列
    layer_types = get_dflash_layer_types(config)
    if layer_types is None:
        # 默认退化为全注意力（非因果编码器模式）
        return -1, AttentionType.ENCODER_ONLY

    if layer_id >= len(layer_types):
        raise ValueError(
            "DFLASH config.layer_types must contain one entry per draft layer. "
            f"Got {len(layer_types)} entries, layer_id={layer_id}."
        )

    layer_type = layer_types[layer_id]
    if layer_type == "full_attention":
        return -1, AttentionType.ENCODER_ONLY
    if layer_type == "sliding_attention":
        sliding_window_size = get_dflash_attention_sliding_window_size(config)
        assert sliding_window_size is not None
        return sliding_window_size, AttentionType.DECODER
    raise ValueError(
        "Unsupported DFLASH draft layer type. "
        f"layer_types[{layer_id}]={layer_type!r}."
    )

class DFlashAttention(nn.Module):
    def __init__(self, config, layer_id: int) -> None:
        super().__init__()
        # ... 其他初始化 ...
        # 原本硬编码 attn_type=AttentionType.ENCODER_ONLY，现在动态获取
        self.sliding_window_size, self.attn_type = _get_dflash_layer_attention_params(
            config, layer_id
        )
        self.attn = RadixAttention(
            num_heads=self.num_heads,
            head_dim=head_dim,
            scaling=self.scaling,
            num_kv_heads=self.num_kv_heads,
            layer_id=layer_id,
            sliding_window_size=self.sliding_window_size,
            attn_type=self.attn_type,
        )

class DFlashDraftModel(nn.Module):
    # ...
    def get_attention_sliding_window_size(self) -> Optional[int]:
        '''对外接口：返回全局滑动窗口大小'''
        return get_dflash_attention_sliding_window_size(self.config)

```

python/sglang/srt/speculative/dflash_utils.py

新增两个关键工具函数，为 dflash.py 和外部查询提供配置读取能力。

```
# python/sglang/srt/speculative/dflash_utils.py
```

```
from collections.abc import Sequence # 新增导入
```

```
def get_dflash_layer_types(config: Any) -> Optional[Sequence[str]]:
    '''从模型配置中读取 layer_types 字段，返回层注意力类型序列。'''
    text_config = _get_text_config(config)
    layer_types = _cfg_get(text_config, "layer_types", _cfg_get(config, "layer_types"))
    if layer_types is None:
        return None
    if isinstance(layer_types, str) or not isinstance(layer_types, Sequence):
        raise ValueError(
            "DFLASH config.layer_types must be a sequence of attention type strings."
        )
    return layer_types

def get_dflash_attention_sliding_window_size(config: Any) -> Optional[int]:
    '''计算 DFlash 滑动窗口大小（转换为 SGLang window_left）。'''
    layer_types = get_dflash_layer_types(config)
    if layer_types is None or "sliding_attention" not in layer_types:
        return None # 没有滑动窗口层时返回 None
    text_config = _get_text_config(config)
    sliding_window = _cfg_get(
        text_config, "sliding_window", _cfg_get(config, "sliding_window")
    )
    if sliding_window is None:
        raise ValueError("DFLASH sliding_attention layers require config.sliding_window.")
    # HuggingFace 风格滑动窗口包含当前 token，SGLang 需要左侧窗口大小
    return int(sliding_window) - 1
```

python/sglang/srt/layers/attention/flashinfer_backend.py

修改滑动窗口更新器和 prefill 元数据生成，使 target-verify 路径能支持 SWA；是注意力后端的核心适配。

```
# python/sglang/srt/layers/attention/flashinfer_backend.py
```

```
class FlashInferIndicesUpdaterPrefill(FlashInferIndicesUpdater):
    def update_sliding_window(self, ...):
        # 新增处理：当 prefix_lens 为 None 时（target-verify 阶段）
        if prefix_lens is None:
            num_accept_tokens = getattr(spec_info, "num_accept_tokens", None)
            if num_accept_tokens is None:
                # 无历史信息时假设 prefix_lens 等于 seq_lens（无前缀）
                prefix_lens = seq_lens
            else:
                # 通过接受 token 数反推 prefix_lens
```

```

        prefix_lens = seq_lens - num_accept_tokens[: seq_lens.shape[0]].to(
            device=seq_lens.device, dtype=seq_lens.dtype
        )
# 后续逻辑使用 prefix_lens 和 sliding_window_size 计算窗口
# ...

def call_begin_forward(self, ...):
    # ... 在 prefill 阶段判断 spec_input_type
    if spec_info.spec_input_type == SpecInputType.DFLASH_VERIFY:
        # DFlash 不使用自定义验证掩码，直接生成 kv_indices 等
        kv_indices, kv_indptr, qo_indptr = (
            spec_info.generate_attn_arg_prefill_with_kv_start_idx(
                req_pool_indices,
                paged_kernel_lens,
                paged_kernel_lens_sum,
                self.req_to_token,
                kv_start_idx, # 传入窗口起始位置，直接裁剪
            )
        )
        custom_mask = None # 没有自定义掩码
    else:
        # 其他 spec 类型（Eagle 等）仍使用原有路径
        kv_indices, kv_indptr, qo_indptr, custom_mask = (
            spec_info.generate_attn_arg_prefill(...)
        )
    # ...

```

评论区精华

1. `accept_length` 变量名：Qiaolin-Yu 指出 `update_sliding_window` 中使用的 `accept_length` 已重命名为 `num_accept_tokens`，dcw02 确认并修复。
 2. `kv_start_idx` 设计：Qiaolin-Yu 询问为何只对 DFlash 传递 `kv_start_idx`（以前 Eagle 路径是否正确），dcw02 解释 DFlash 不使用自定义验证掩码，可直接裁剪；Eagle 使用自定义掩码需要掩码列对齐，不能直接套用。
 3. 代码复用：Qiaolin-Yu 建议将 `prefix_lens` 为 `None` 的推导移到 `init` 函数，dcw02 改为复用 `self.sliding_window_size`。
- `accept_length` 重命名为 `num_accept_tokens (correctness)`: dcw02 确认并修复
 - `kv_start_idx` 仅用于 DFlash 的设计 (design): dcw02 解释 DFlash 不使用自定义验证掩码，可直接裁剪 KV 索引；Eagle 使用自定义掩码需要掩码列对齐，不适用
 - 建议将 `prefix_lens` 推导移到 `init` 函数 (design): dcw02 改为复用 `self.sliding_window_size`

风险与影响

- 风险：主要风险在于修改了 FlashInfer 后端的 `update_sliding_window` 通用路径，可能影响现有滑动窗口场景（如 Eagle 或非 DFlash 验证）。但改动仅针对 `prefix_lens` 为 `None`

时添加回退逻辑，且专门分支 DFLASH_VERIFY 才使用 kv_start_idx，Eagle 等路径仍沿用自定义掩码路径，回归风险可控。测试仅覆盖 FlashInfer 后端，其他后端（如 triton）的 DFlash SWA 尚未验证。

- 影响：用户侧：DFlash 草稿模型现在可以配置 layer_types 使用滑动窗口层，降低 KV 缓存开销。系统侧：注意力后端需要为新的 spec_input_type 分支提供支持；后续增加非 FlashInfer 后端的 SWA verify 支持可基于此范式扩展。团队侧：需维护注意力类型配置契约，但整体变更有限。
- 风险标记：核心路径变更，测试覆盖仅限 FlashInfer 后端，缺少对非 FlashInfer 后端的验证

关联脉络

- PR #23000 DFlash layer structure / scheduler integration (原始大 PR): 此 PR 是从 #23000 拆分的子任务，专注于滑动窗口注意力层支持。