

PR #30111 完整报告

sgl-project/sglang

[Fix] Fix DSA indexer fusion for NeoX RoPE

合并时间: 2026-07-04 18:20

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/30111>

执行摘要

- 一句话: 修复 NeoX RoPE 下 DSA indexer fusion 精度回退
- 推荐动作: 本 PR 值得阅读, 展示了如何将全局开关细化为实例级开关以解决模型兼容性问题。代码改动简洁、测试阈值提升合理。

功能与动机

PR Body 指出根因: DSA indexer fusion 仅由全局环境变量控制, NeoX-style RoPE 模型仍会进入 fusion 路径, 但该路径不兼容 `is_neox_style=True` 的旋转处理, 导致精度回退。

实现拆解

1. 移除模块级全局开关: 在 `dsa_indexer.py` 中删除 `_use_dsa_indexer_fusion` 模块级变量, 不再依赖全局布尔值。
2. 添加实例级 fusion 属性: 在 `DSAIndexer.__init__` 中新增 `self.use_dsa_indexer_fusion`, 其值为 `_is_cuda` and not `SGLANG_DISABLE_DSA_INDEXER_FUSION` and not `is_neox_style`。这样每个 `indexer` 实例根据自身 `is_neox_style` 参数决定是否启用 fusion。
3. 替换所有引用: 将原本 `_use_dsa_indexer_fusion` 的 7 处引用全部改为 `self.use_dsa_indexer_fusion`, 包括构造函数分支、`_maybe_rotate`、`_get_q_k_bf16`、`_forward_cuda_k_only`、`forward_cuda` 等。
4. 恢复全局默认启用: 在 `environ.py` 中将 `SGLANG_DISABLE_DSA_INDEXER_FUSION` 从 `EnvBool(True)` 改回 `EnvBool(False)`, 保持默认 fusion 开启。
5. 调整测试阈值: 在 `test_deepseek_v32_indexcache.py` 中将 GSM8K 准确率下限从 0.93 提升至 0.935, 反映修复后的更高精度。

关键文件:

- `python/sglang/srt/layers/attention/dsa/dsa_indexer.py` (模块 注意力; 类别 source; 类型 core-logic): 核心逻辑变更: 移除模块级 `_use_dsa_indexer_fusion`, 新增实例级 `self.use_dsa_indexer_fusion`, 并替换所有引用。
- `python/sglang/srt/environ.py` (模块 环境配置; 类别 source; 类型 core-logic): 恢复 `SGLANG_DISABLE_DSA_INDEXER_FUSION` 默认值为 `False`, 使 fusion 默认开启。
- `test/registered/8-gpu-models/test_deepseek_v32_indexcache.py` (模块 测试; 类别 test; 类型 test-coverage): 提升 GSM8K 测试阈值, 反映修复后的更高精度。

关键符号: init, _maybe_rotate, _get_q_k_bf16, _forward_cuda_k_only, forward_cuda

关键源码片段

python/sglang/srt/layers/attention/dsa/dsa_indexer.py

核心逻辑变更: 移除模块级 `_use_dsa_indexer_fusion`, 新增实例级 `self.use_dsa_indexer_fusion`, 并替换所有引用。

```
# dsa_indexer.py
```

```
class DSAIndexer(nn.Module):
    def __init__(
        self,
        hidden_size: int,
        index_n_heads: int,
        index_head_dim: int,
        rope_head_dim: int,
        index_topk: int,
        q_lora_rank: int,
        max_position_embeddings: int,
        rope_theta: float,
        layer_id: int,
        scale_fmt: Optional[str],
        block_size: int = 128,
        rope_scaling: Optional[Dict[str, Any]] = None,
        is_neox_style: bool = True,
        prefix: str = "",
        quant_config: Optional[QuantizationConfig] = None,
        alt_stream: Optional[torch.cuda.Stream] = None,
    ):
        super().__init__()
        # ... 其他初始化 ...
        # 实例级 fusion 标志: 仅在 CUDA + 未禁用 + 非 NeoX-style 时启用
        self.use_dsa_indexer_fusion = (
            _is_cuda
            and not envs.SGLANG_DISABLE_DSA_INDEXER_FUSION.get()
            and not is_neox_style
        )
        # ...

        # 根据 fusion 标志选择不同的线性层配置
        if self.use_dsa_indexer_fusion:
            self.wk_weights_proj = ReplicatedLinear(
                self.hidden_size,
                self.head_dim + self.n_heads,
                bias=False,
                params_dtype=torch.bfloat16,
                prefix=add_prefix("wk_weights_proj", prefix),
            )
```

```

else:
    self.wk = ReplicatedLinear(
        self.hidden_size,
        self.head_dim,
        bias=False,
        quant_config=quant_config,
        prefix=add_prefix("wk", prefix),
    )
    self.weights_proj = ReplicatedLinear(
        self.hidden_size,
        self.n_heads,
        bias=False,
        quant_config=quant_config,
        prefix=add_prefix("weights_proj", prefix),
    )

def _maybe_rotate(self, x: torch.Tensor) -> torch.Tensor:
    # Fusion 路径下跳过 Hadamard 旋转以保持与 decode 读取的一致性
    return x if self.use_dsa_indexer_fusion else rotate_activation(x)

def _get_q_k_bf16(self, x, forward_batch, num_tokens=None):
    # 多个方法中引用实例级标志
    if self.use_dsa_indexer_fusion:
        key, weights_raw = self._fused_k_weights(x)
    else:
        key, _ = self.wk(x)
        weights_raw, _ = self.weights_proj(x)
    # ...

def forward_cuda(self, ...):
    # 在 forward 中根据 fusion 标志决定是否短路 weights_proj LoRA 处理
    weights_proj_lora = not self.use_dsa_indexer_fusion and getattr(
        self.weights_proj, "set_lora", False
    )
    if (
        self.use_dsa_indexer_fusion
        and not in_pieewise_or_breakable_cuda_graph
        and forward_batch.attn_cp_metadata is None
    ):
        # 使用融合路径
        ...

```

[test/registered/8-gpu-models/test_deepseek_v32_indexcache.py](#)

提升 GSM8K 测试阈值，反映修复后的更高精度。

```

# test_deepseek_v32_indexcache.py (片段)
if is_in_ci():
    write_github_step_summary(
        f"### test_gsm8k (deepseek-v32)\n" f'{metrics["accuracy"]:.3f}\n'

```

```
)  
self.assertGreater(metrics["accuracy"], 0.935) # 从 0.93 提升至 0.935
```

评论区精华

无人工 review 评论；仅 Gemini Code Assist 自动回复无实质反馈。

- 暂无高价值评论线程

风险与影响

- 风险：低风险。变更集中在 1 个核心文件的 7 处引用替换和 1 个环境变量默认值、1 个测试阈值调整。逻辑清晰：仅当 `is_neox_style=False` 时启用 fusion，不影响其他模型。可能风险：self.use_dsa_indexer_fusion 在 _maybe_rotate 中控制是否跳过 Hadamard 旋转，若后续增加新的旋转风格可能需扩展该条件。
- 影响：影响 DeepSeek V3.2 及所有使用 NeoX-style RoPE 的 DSA 模型（如 GLM5.2 MHA）。修复后这些模型可正常使用 DSA indexer fusion（默认开启），避免精度回退。GLM5.2 相关测试 [#29959] 已覆盖 MHA 场景。
- 风险标记：核心路径变更，测试阈值调整

关联脉络

- PR #29959 [DSA][GLM5.2] Index Share for MHA: 同一 DSA indexer 模块的关联 PR，涉及 MHA 的 indexer 共享跳过逻辑，可能与本修复的 NeoX 兼容性问题有关。
- PR #29843 [trtllm_mha] Fuse cuda-graph metadata rebuild into one triton kernel: 同为注意力后端优化，涉及 CUDA graph 融合，与本 PR 的 indexer fusion 修复无直接关联但属于同类优化方向。