

PR #33616 完整报告

sgl-project/sglang

feat: Add flashinfer mHC fusion for DSV4

合并时间: 2026-08-07 16:01

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/33616>

执行摘要

- 一句话: 新增 DSV4 mHC 的 FlashInfer 后端, 默认关闭, 可切换备选融合实现。
- 推荐动作: 该 PR 代码量小 (86 行), 但值得精读 `_flashinfer_mhc_pre_num_splits` 的 split-K 选择策略与 `hc_pre/hc_post` 的分支优先级安排, 作为多后端融合实现的参考。若计划在生产环境启用, 建议先补充数值一致性测试, 并关注 FlashInfer 上游 API 变更。

功能与动机

PR body 明确说明: 为 DSV4 的 mHC 增加 flashinfer backend, 当前 tilelang mhc fusion 仍更快, 但当 tilelang 不可用时 flashinfer 可作为 performant alternative, 且预期 flashinfer 很快支持 full pre fusion 和 post+pre fusion。目的是在保持默认路径不变的前提下, 为不同部署环境提供灵活的高性能选择。

实现拆解

1. 新增环境变量开关: 在 `python/sglang/srt/envron.py` 的 CUDA kernels 配置区新增 `SGLANG_OPT_USE_FLASHINFER_MHC = EnvBool(False)`, 默认关闭, 与既有 `TileLang/Aiter` 开关并列, 确保不改变默认行为。
2. 新增 FlashInfer pre 融合辅助函数: 在 `python/sglang/srt/models/deepseek_v4.py` 模块顶层定义 `_FLASHINFER_MHC_PRE_SPLITS` 常量 (允许的 split-K 取值)、缓存 SM 数的 `_cuda_sm_count()`、按 token 数与 hidden size 自动选择 split-K 的 `_flashinfer_mhc_pre_num_splits()`, 以及执行 pre 融合的 `_flashinfer_hc_pre()`。
`_flashinfer_hc_pre` 先调用 `deepgemm` 封装 `tf32_hc_prenorm_gemm` 计算 `dot_mix` 与 `sqrsum` (支持 split-K), 再调用 `flashinfer.mhc.mhc_pre_big_fuse` 完成 sinkhorn 归一化与组合矩阵计算, 返回 `layer_input, post, comb`。
3. 在 `hc_pre` 与 `hc_post` 方法中插入 FlashInfer 分支: 在 `hc_pre` 中, 将 `SGLANG_OPT_USE_FLASHINFER_MHC` 分支放在 `TileLang` 分支之前, 调用 `_flashinfer_hc_pre` 并返回四元组 `(y, post, comb, False)`; 在 `hc_post` 中, 同样优先走 `flashinfer.mhc.mhc_post`。分支顺序保证开启该 flag 时完全绕开 `TileLang` 与 `Aiter` 路径。
4. 测试与验证配套: 本 PR 未新增单元测试, 仅作者在 body 中提供 GSM8K 手工验证结果 (accuracy 0.975), 并使用 `--moe-runner-backend flashinfer_mxfp4` 等参数复现。CI 仅触发既有测试流程, 未针对新路径增加覆盖。

关键文件:

- python/sglang/srt/models/deepseek_v4.py (模块 模型实现; 类别 source; 类型 core-logic; 符号 _cuda_sm_count, _flashinfer_mhc_pre_num_splits, _flashinfer_hc_pre) : 核心变更文件: 新增 FlashInfer mHC pre/post 融合辅助函数和分支, 是 DSV4 模型计算路径的具体实现。
- python/sglang/srt/envron.py (模块 环境变量; 类别 source; 类型 configuration) : 新增 SGLANG_OPT_USE_FLASHINFER_MHC 环境变量开关, 默认关闭, 是启用该后端的入口。

关键符号: _flashinfer_hc_pre, _flashinfer_mhc_pre_num_splits, _cuda_sm_count, hc_pre, hc_post

关键源码片段

python/sglang/srt/models/deepseek_v4.py

核心变更文件: 新增 FlashInfer mHC pre/post 融合辅助函数和分支, 是 DSV4 模型计算路径的具体实现。

```
# FlashInfer 的 mhc_pre_big_fuse 仅接受这些 split-K 取值
_FLASHINFER_MHC_PRE_SPLITS = (1, 2, 4, 8, 16)
```

```
@functools.cache
```

```
def _cuda_sm_count() -> int:
```

```
    # 缓存 SM 数量, 避免反复查询设备属性
```

```
    return torch.cuda.get_device_properties(0).multi_processor_count
```

```
def _flashinfer_mhc_pre_num_splits(num_tokens: int, hc_hidden_size: int) -> int:
```

```
    # 根据 token 数与 hidden size 估算网格大小, 再结合 SM 数挑选合理 split-K
```

```
    block_m = block_k = 64
```

```
    grid_m = (num_tokens + block_m - 1) // block_m
```

```
    num_block_k = (hc_hidden_size + block_k - 1) // block_k
```

```
    raw = max(1, min(_cuda_sm_count() // max(grid_m, 1), num_block_k // 4))
```

```
    best = 1
```

```
    for split in _FLASHINFER_MHC_PRE_SPLITS:
```

```
        if split <= raw:
```

```
            best = split
```

```
    return best
```

```
def _flashinfer_hc_pre(
```

```
    x: torch.Tensor,
```

```
    hc_fn: torch.Tensor,
```

```
    hc_scale: torch.Tensor,
```

```
    hc_base: torch.Tensor,
```

```
    *,
```

```
    rms_eps: float,
```

```
    hc_eps: float,
```

```
    sinkhorn_iters: int,
```

```

) -> Tuple[torch.Tensor, torch.Tensor, torch.Tensor]:
    from flashinfer.mhc import mhc_pre_big_fuse
    from sglang.srt.layers.deep_gemm_wrapper.entrypoint import tf32_hc_prenorm_gemm

    num_tokens, hc_mult, hidden_size = x.shape
    hc_hidden_size = hc_mult * hidden_size
    mix_dim = hc_fn.shape[0] # hc_mult * (2 + hc_mult) == 24
    n_splits = _flashinfer_mhc_pre_num_splits(num_tokens, hc_hidden_size)

    # 先用 deepgemm 计算 pre-norm 的 dot_mix 与平方和，再交给 flashinfer 融合
    dot_mix = torch.empty(
        (n_splits, num_tokens, mix_dim), dtype=torch.float32, device=x.device
    )
    sqrsum = torch.empty((n_splits, num_tokens), dtype=torch.float32, device=x.device)
    tf32_hc_prenorm_gemm(
        x.reshape(num_tokens, hc_hidden_size), hc_fn, dot_mix, sqrsum, n_splits
    )
    if n_splits == 1:
        dot_mix = dot_mix.squeeze(0)
        sqrsum = sqrsum.squeeze(0)

    post, comb, layer_input = mhc_pre_big_fuse(
        dot_mix,
        sqrsum,
        x,
        hc_scale,
        hc_base,
        hc_hidden_size,
        rms_eps=rms_eps,
        mhc_pre_eps=hc_eps,
        mhc_sinkhorn_eps=hc_eps,
        mhc_post_mult_value=_MHC_POST_MULT_VALUE,
        sinkhorn_repeat=sinkhorn_iters,
        num_splits=n_splits,
    )
    return layer_input, post.squeeze(-1), comb

```

评论区精华

该 PR 的 review 讨论很少：维护者 b8zhong 直接批准（APPROVED），Issue 中仅有一条 [/rerun-failed-ci](#) 请求，没有实质性的技术争论。从 PR body 可见作者已明确说明当前 TileLang 更快、FlashInfer 是备选方案，因此没有出现设计分歧。

- 暂无高价值评论线程

风险与影响

- 风险：

1. 数值一致性风险: `hc_pre/hc_post` 是 DSV4 核心计算路径, 新分支默认关闭, 但一旦启用则所有层均走 FlashInfer 实现, 与 TileLang/torch 参考实现可能产生数值差异, 作者仅用 GSM8K 抽查, 未提供逐层对齐测试。
2. split-K 启发式风险: `_flashinfer_mhc_pre_num_splits` 依赖 `torch.cuda.get_device_properties(0).multi_processor_count`, 在多卡环境中默认取 0 号设备, 若各卡 SM 数不一致 (现实中极少) 可能选错 split 数; 且该启发式未覆盖极端形状 (如 `num_tokens` 极小或 `hc_hidden_size` 极大)。
3. 外部库 API 兼容性: 直接依赖 `flashinfer.mhc.mhc_pre_big_fuse` 与 `flashinfer.mhc.mhc_post`, FlashInfer 版本升级可能改变签名、数值行为或引入回归, 且当前无版本约束或降级处理。
4. 测试覆盖缺失: 没有针对新分支的单元测试或数值对齐测试, 未来重构 TileLang 路径时可能无意破坏该分支而不被 CI 发现。
 - 影响: 影响范围限定在 DSV4 模型, 影响程度低 (默认关闭)。用户需要显式设置 `SGLANG_OPT_USE_FLASHINFER_MHC=1` 且安装包含 `flashinfer.mhc` 的版本才会启用; 未设置环境变量的现有部署完全不受影响。对团队而言, 该 PR 建立了 mHC 多后端融合的雏形, 为后续 FlashInfer 全融合能力接入铺路, 同时展示了 split-K 自动选择这一可复用设计。
 - 风险标记: 缺少测试覆盖, 默认关闭降低风险, 依赖外部库新 API, 数值一致性未自动化验证

关联脉络

- 暂无明显关联 PR