

VWN_V2 fuse kernel接入big_op

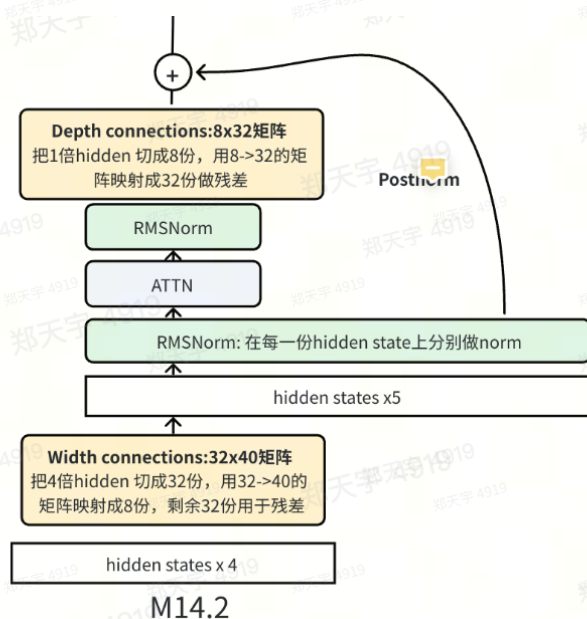
1.v2的一些改进

1. 分块处理

当前 $n=32$, $m=8$, 计算时, A和B矩阵拼在一起, 是 $[:,mn,m,:]$, $mn=48$, 直接计算从48 padding到64比较浪费, 因此分成32和16来处理

2. width_connection和depth_connection都做fuse kernel

3. residual放在depth_connection计算并相加



如果是postnorm, 这里需要对每一份hidden state做norm, 而这个操作可以fuse到depth connections, 且对于 $5 \times h$ 里的其中的后面 $4 \times h$, 可以不在width connection中计算, 而是放到depth connection计算, 并和resnorm做fuse。

4. width connections仅输出h大小的数据用于attn/ffn计算, 并且beta变成 alpha_res_beta, alpha_res即用于算residual那部分的alpha。

2.v2目前可优化的点

1. Todo: width connection的输出一定会做resnorm, 因此也可以fuse进kernel里面

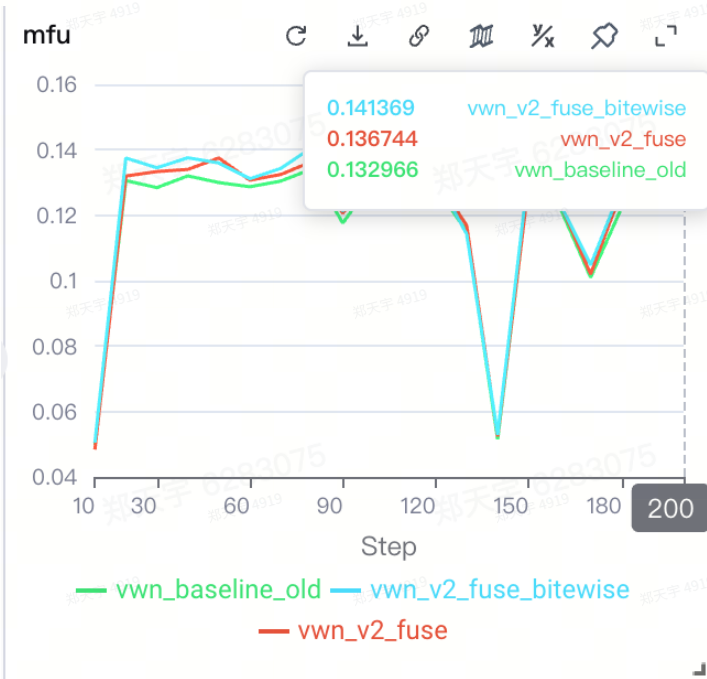
3.当前的实验效果

实验的trail:https://seed.byteintl.net/experiment/tracking/detail?Id=project_20260317_6645be31&selectedTrial=&tab=CHART&viewId=view_20260508_c543a772

loss：不能bitwise对齐，误差在预期范围内，符合预期。fuse_v2 验证是deterministic的。



mfu：
mfu提升0.009符合预期，且多次实验均显著高于baseline

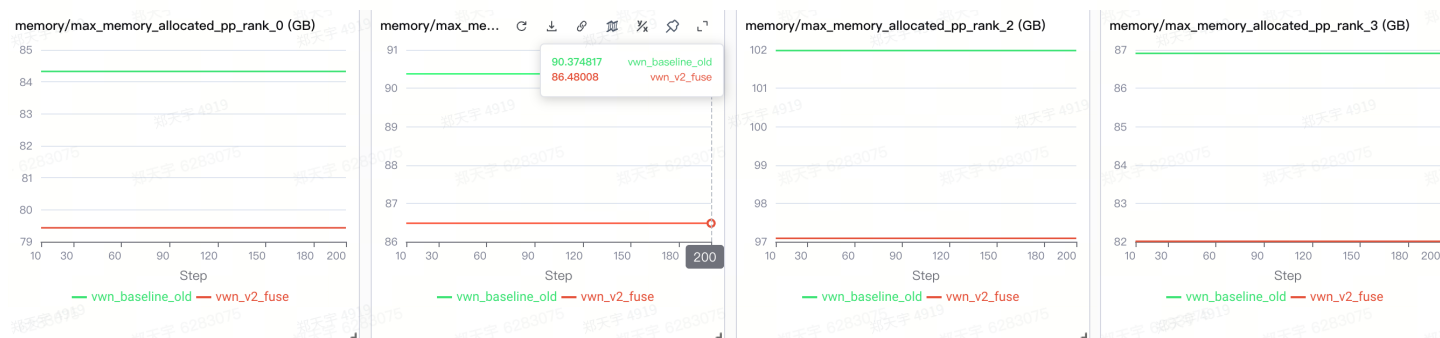


二次实验，mfu提升比较稳定



memory:

memory在pp-3上节省约4个G，符合预期，且每个pp rank上均显著下降

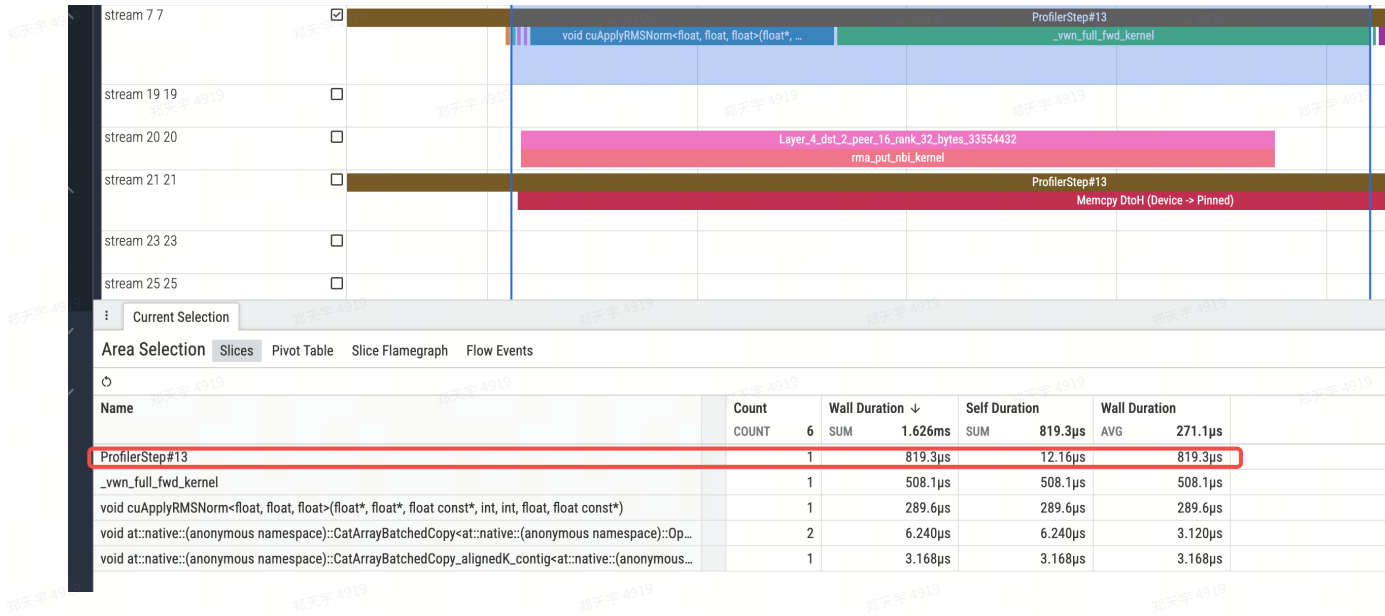


kernel执行时间:

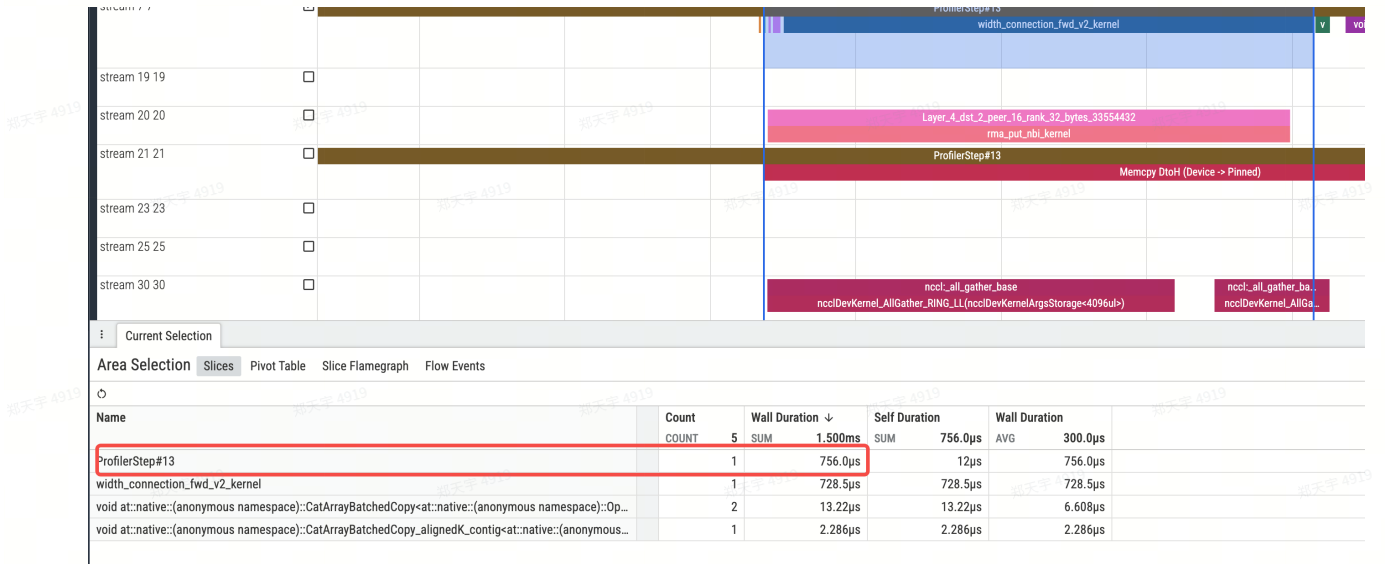
- width_connection

819us->756us

before:



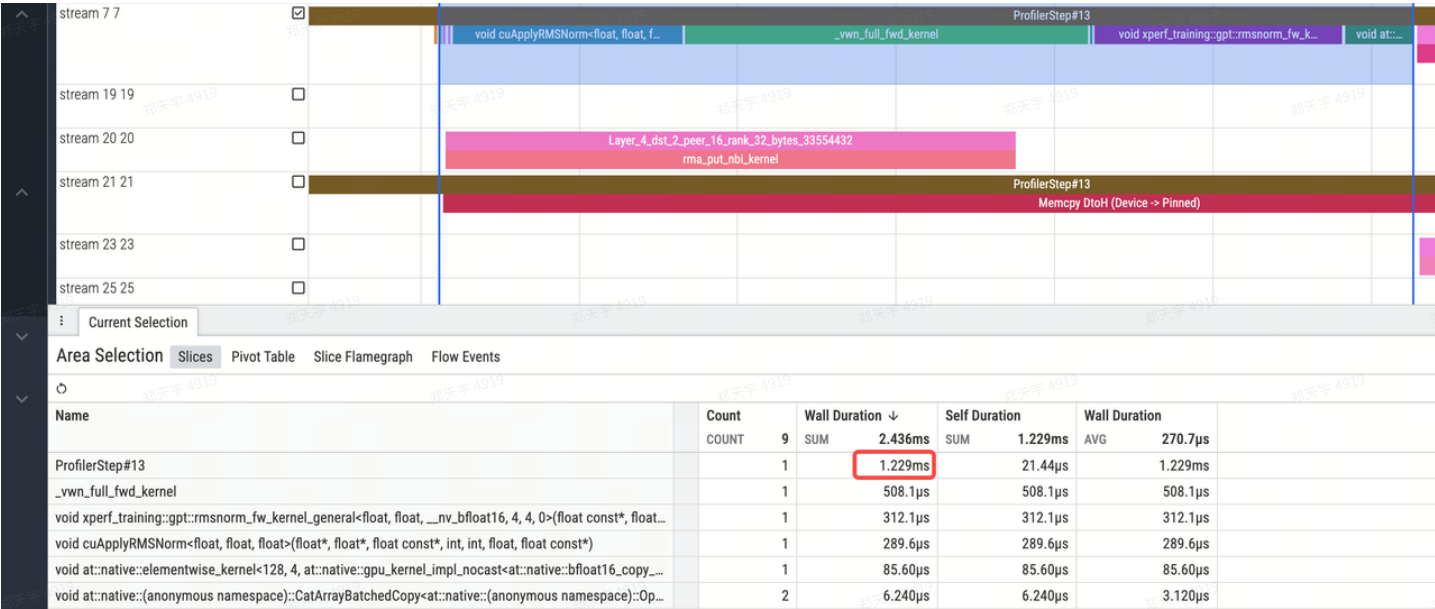
after:



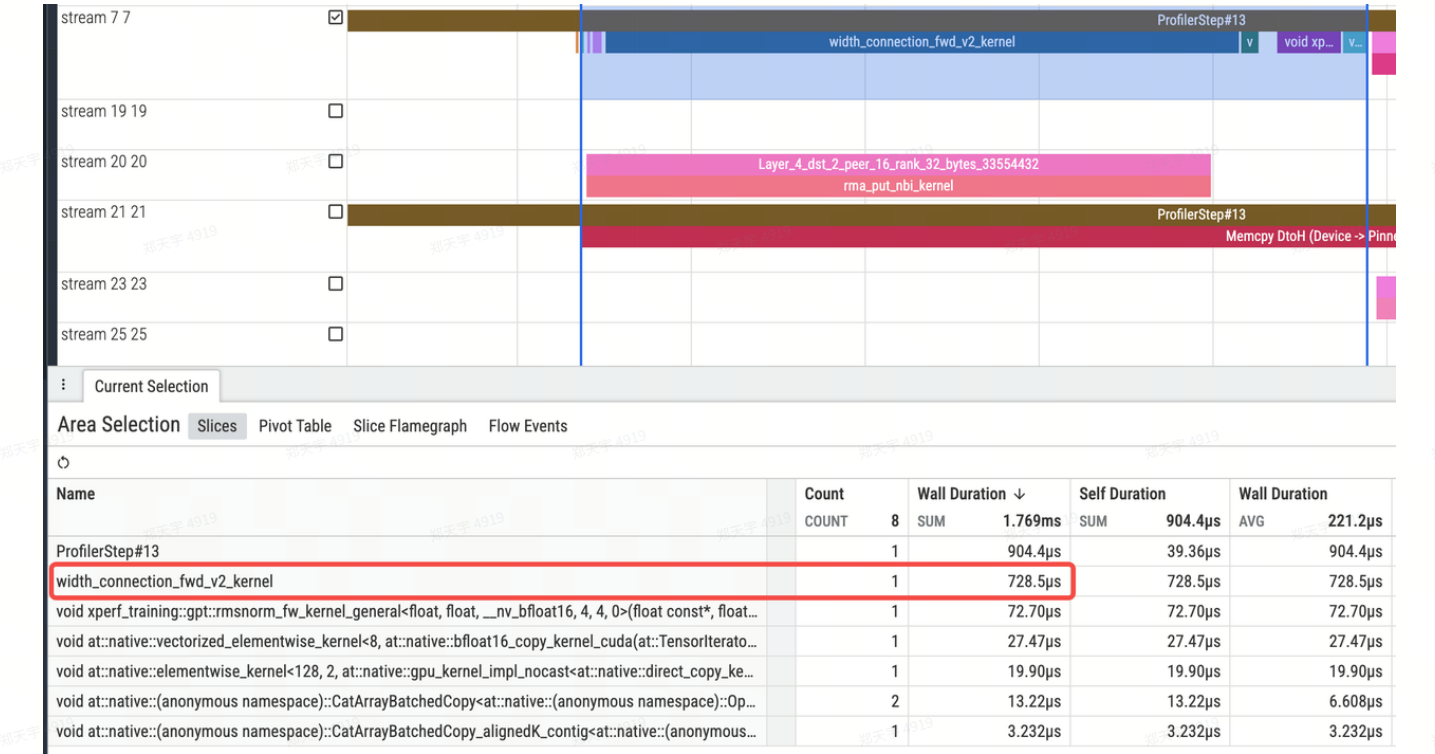
- width_connection+rmsnorm

1.229->0.904 ms

before:



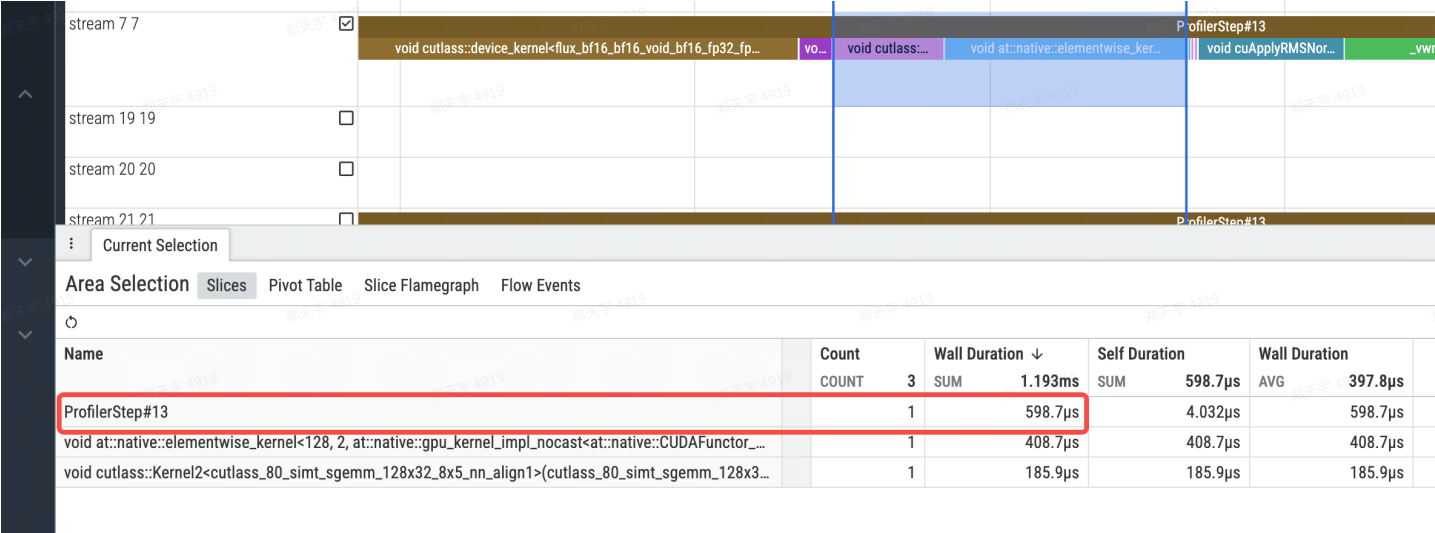
after:



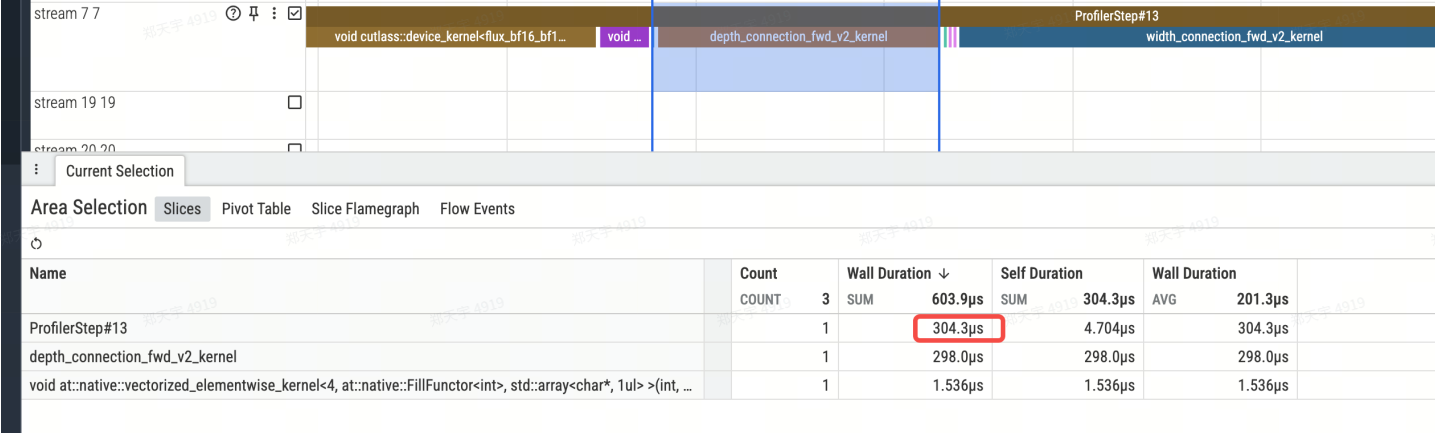
- depth_connection

598us->304us

before:



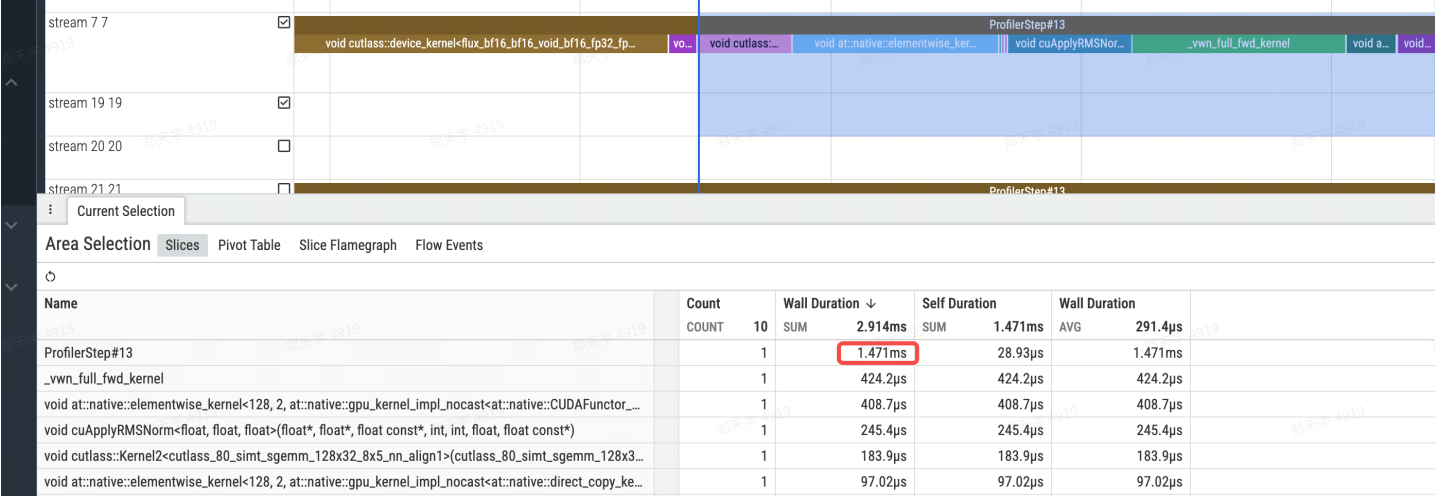
after:



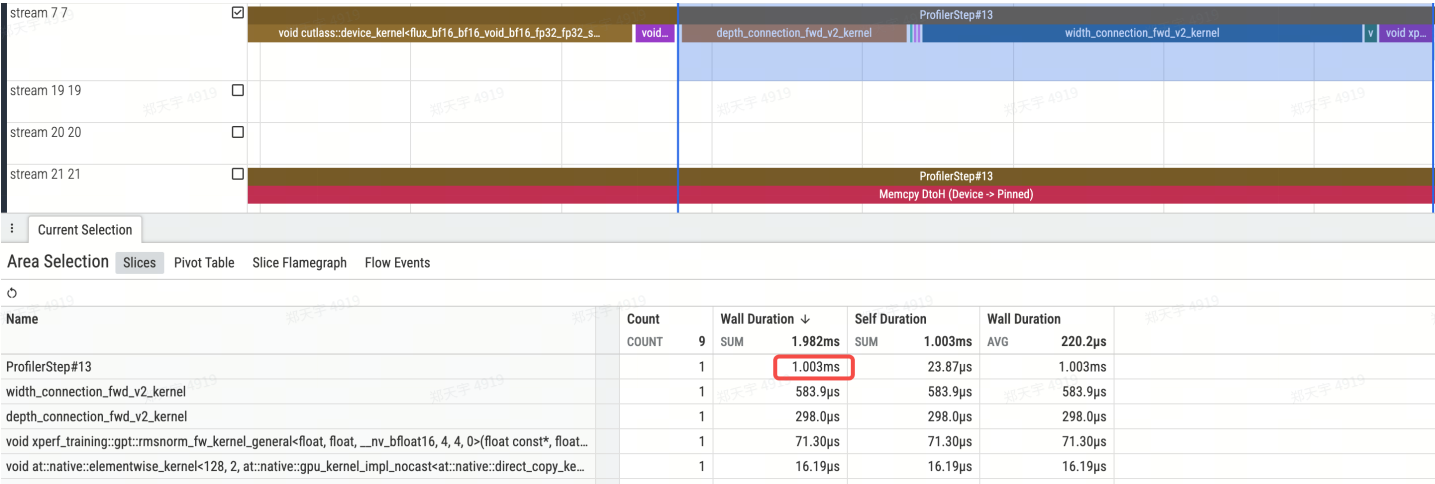
- 第一个depth_connection+第二个width_connection+rmsnorm

1.47ms->1ms

before:



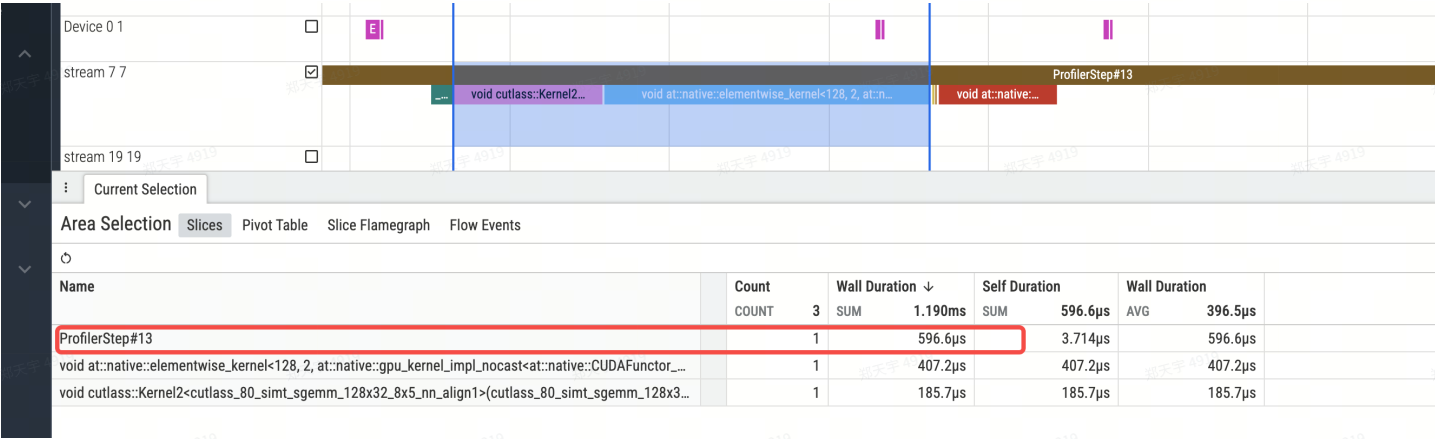
after:



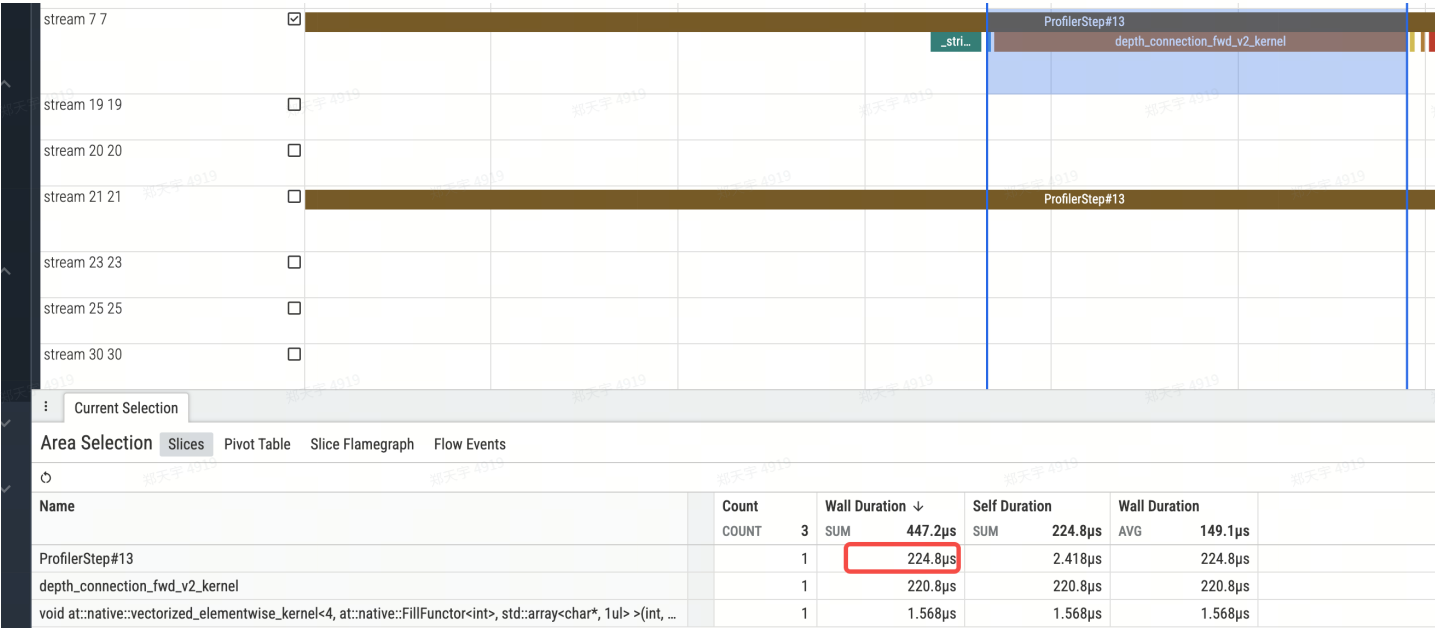
• 第二个depth_connection

596us->224us

before:



after:



vwn_kernel更新后大约节省1.2ms

4.重计算



depth_connection_bwd_kerenl在执行前，需要先执行2个width_connection_fwd以及一个depth_connection。

链式法则比较好理解，并且从参数上看：depth_connection_bwd2需要beta2，而beta2由width_connection_fwd2计算得到，而width_connection_fwd2的输入即经过残差连接的4h输入，由depth_connection_fwd1得到，而depth_connection_fwd1又需要width_connection_fwd1的beta1，所以必须要做这些计算。

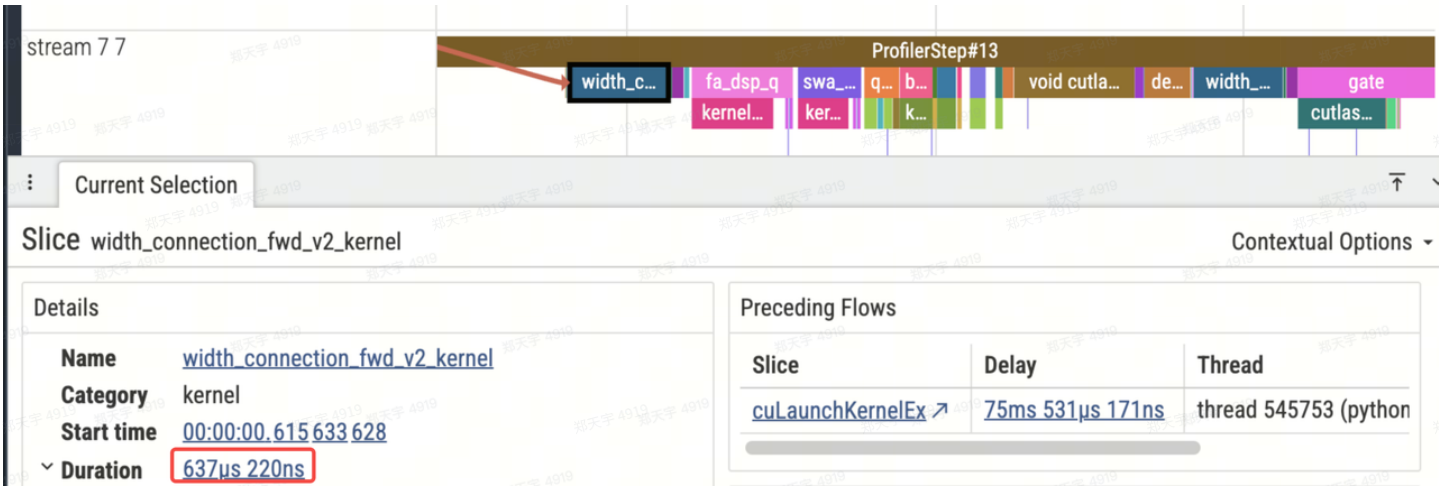
5.记录一下v2的使用(同时记录 triton 3.5版 deterministic的实验)

为了满足m=8的矩阵乘，得用triton>=3.5以上的版本，之前的版本都不支持，最低shape m, n, k都为16。

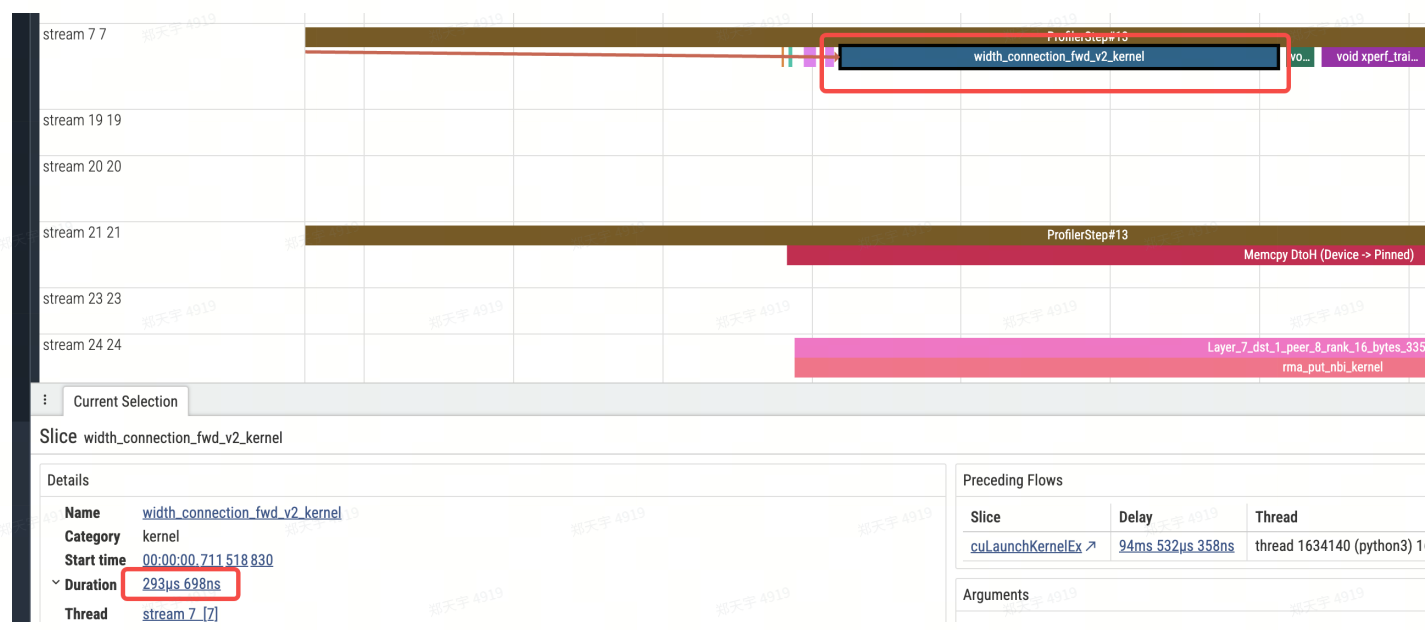
```
data/aml/mlsys/triton_dist_pt29_cu131 1.0.0.1 /opt/tiger/triton_dist /opt/tiger/triton_dist/python
_tt35
```

SCM链接：<https://cloud.bytedance.net/scm/detail/532532/versions?x-resource-account=public&x-bc-region-id=bytedance>

如果padding到16，width_connection的时间会增加：



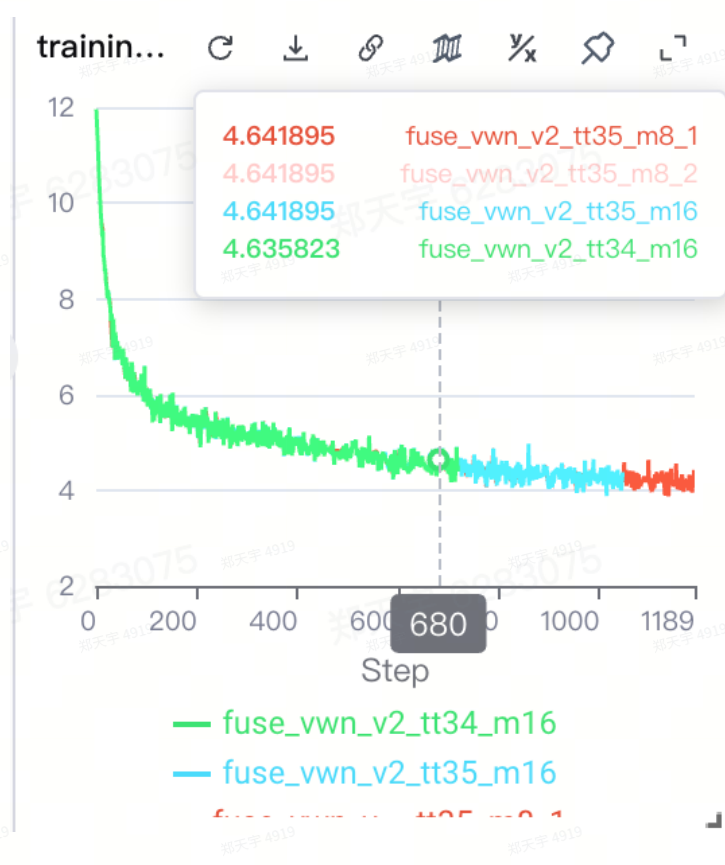
更换triton后，保持BLOCK_M=8，则width_connection的时间缩减到原来的1/2，符合预期：



- triton3.5测试:

测试链接：https://seed.byteintl.net/experiment/tracking/detail?Id=project_20260317_6645be31&selectedTrial=&tab=CHART&viewId=view_20260510_6e15bdce

loss：符合预期，blockm=8/16都是bitwise对齐的



mfu: 看起来差不多



memory: 峰值显存不变

