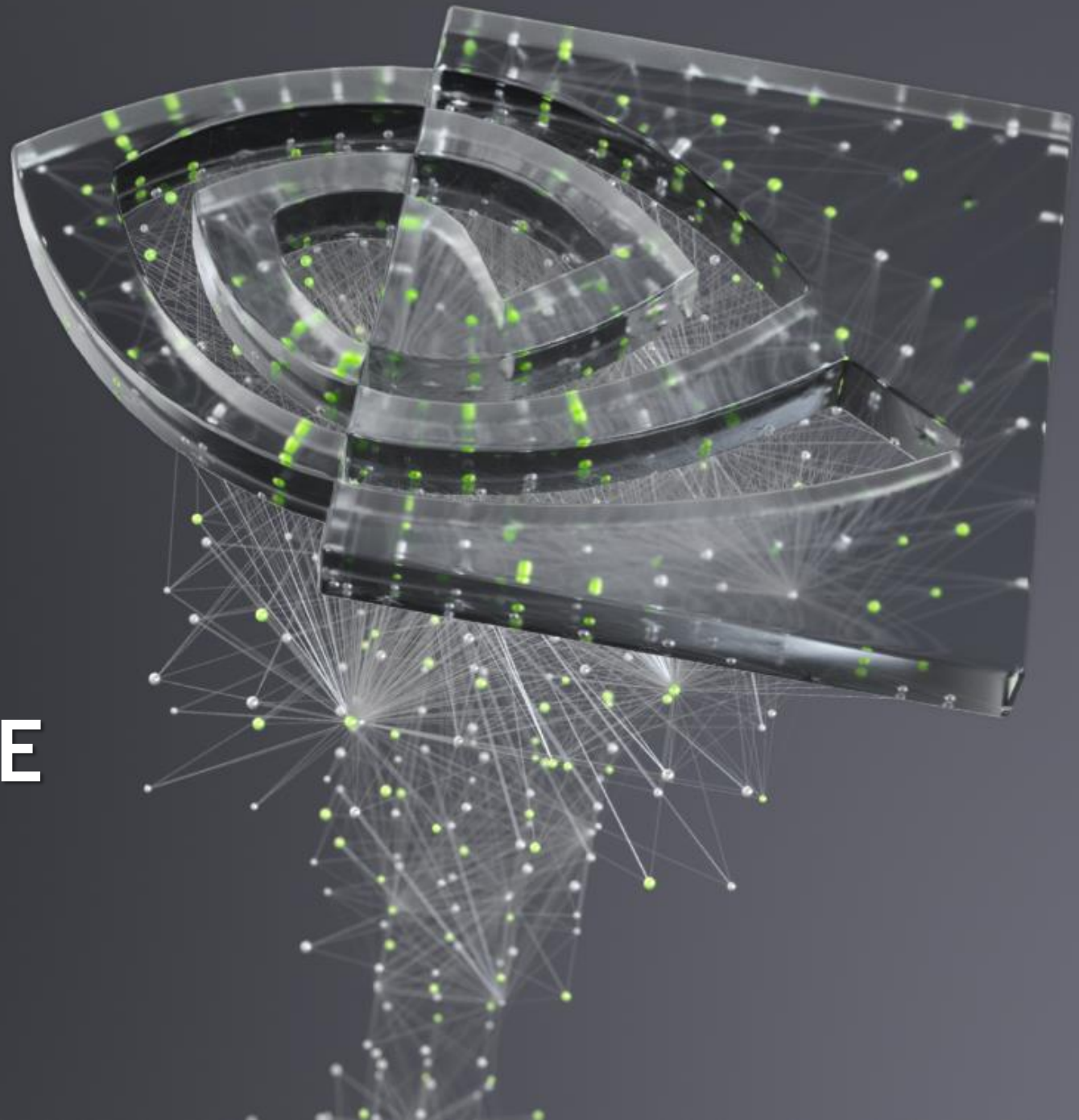




CUDA PERFORMANCE OPTIMIZATION

Markus Hrywniak, DevTech Compute, April 29, 2021



OUTLINE

What you will (not) learn this session

Memory access patterns

- Coalesced access

Branch divergence

Using profiling tools (some more) for Kernel level

Also important, but not covered here:

- Exposing enough parallelism

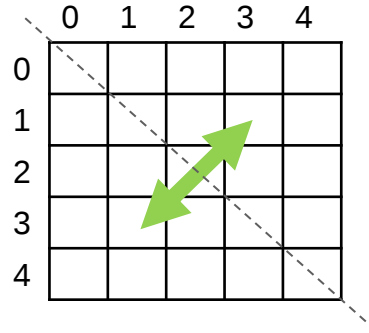
- Expressing data locality

- Application-wide profiling

MOTIVATING EXAMPLE: MATRIX TRANSPOSE

CPU VERSION:

```
void transpose(
Integer* a_trans,
Integer* a,
Integer n)
{
    for (Integer row = 0; row < n; ++row)
        for (Integer col = 0; col < n; ++col)
            a_trans[col][row] = a[row][col];
}
```



GPU VERSION:

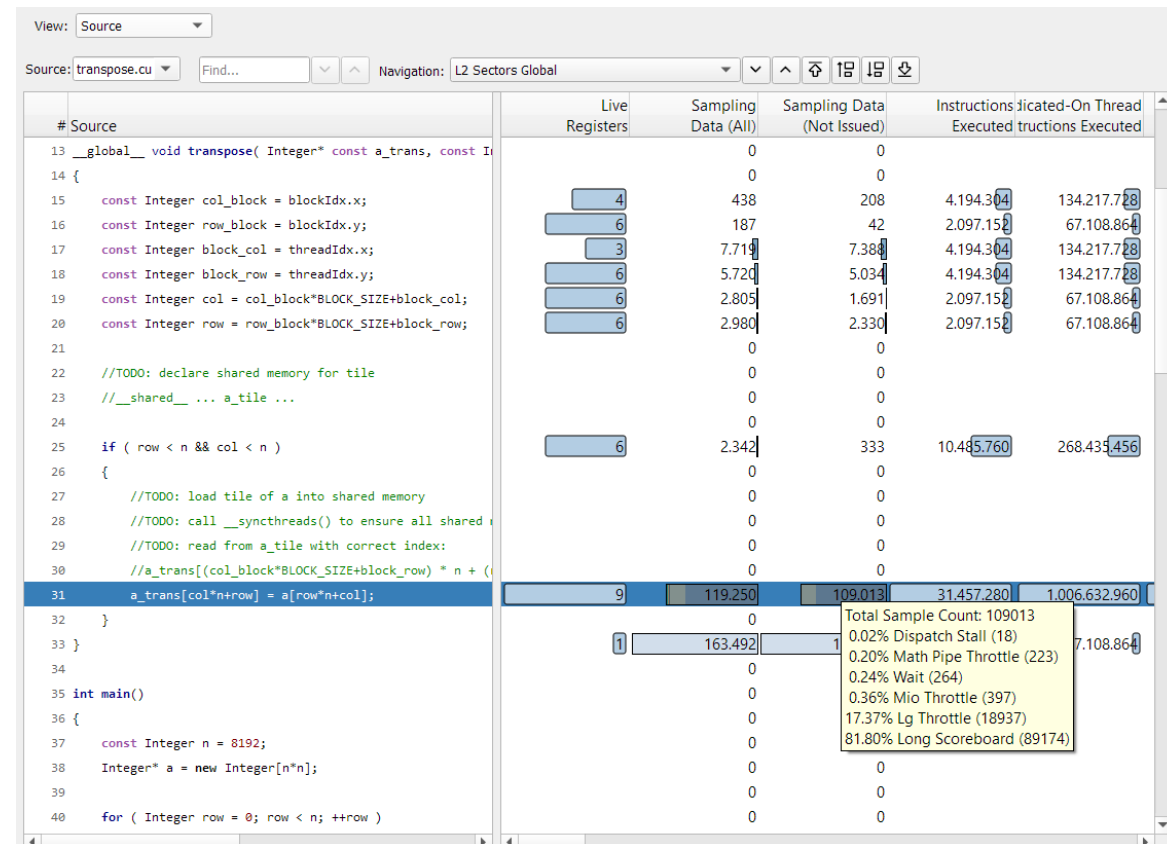
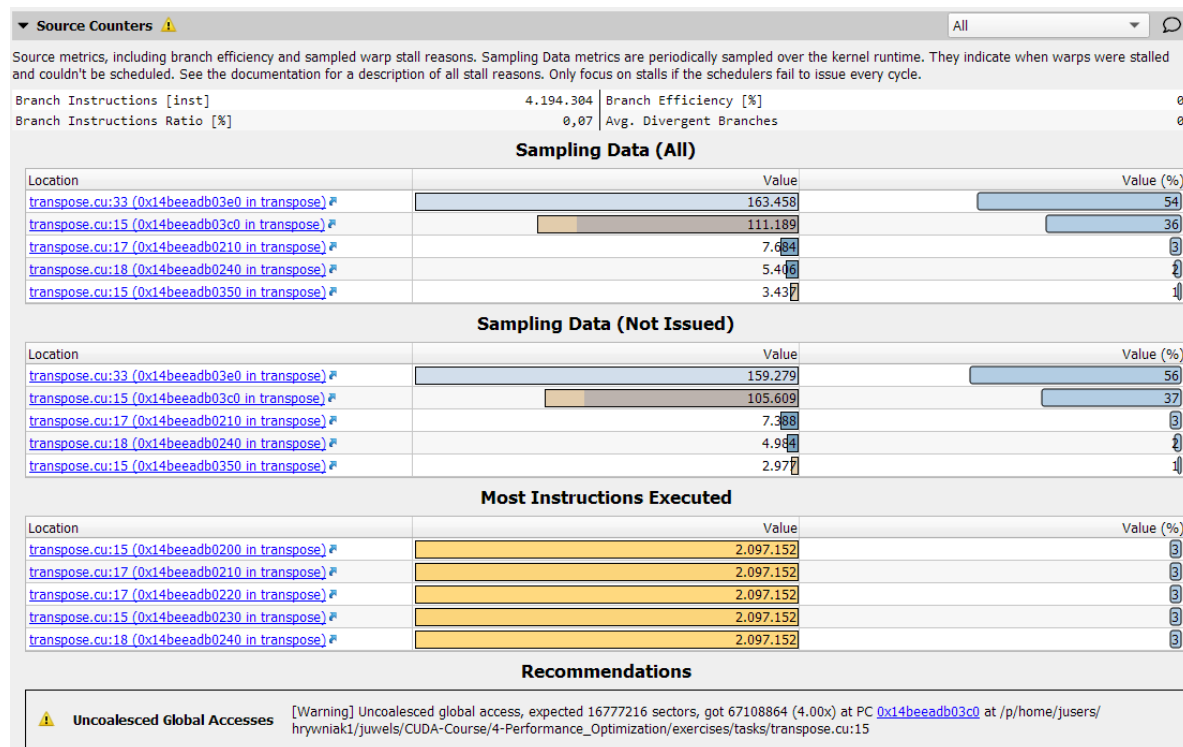
```
__global__ void transpose(
Integer* a_trans,
Integer* a,
Integer n)
{
    Integer row =
        blockIdx.y*blockDim.y+threadIdx.y;
    Integer col =
        blockIdx.x*blockDim.x+threadIdx.x;

    if (row < n && col < n)
        a_trans[col][row] = a[row][col];
}
```



Row-major
ordering

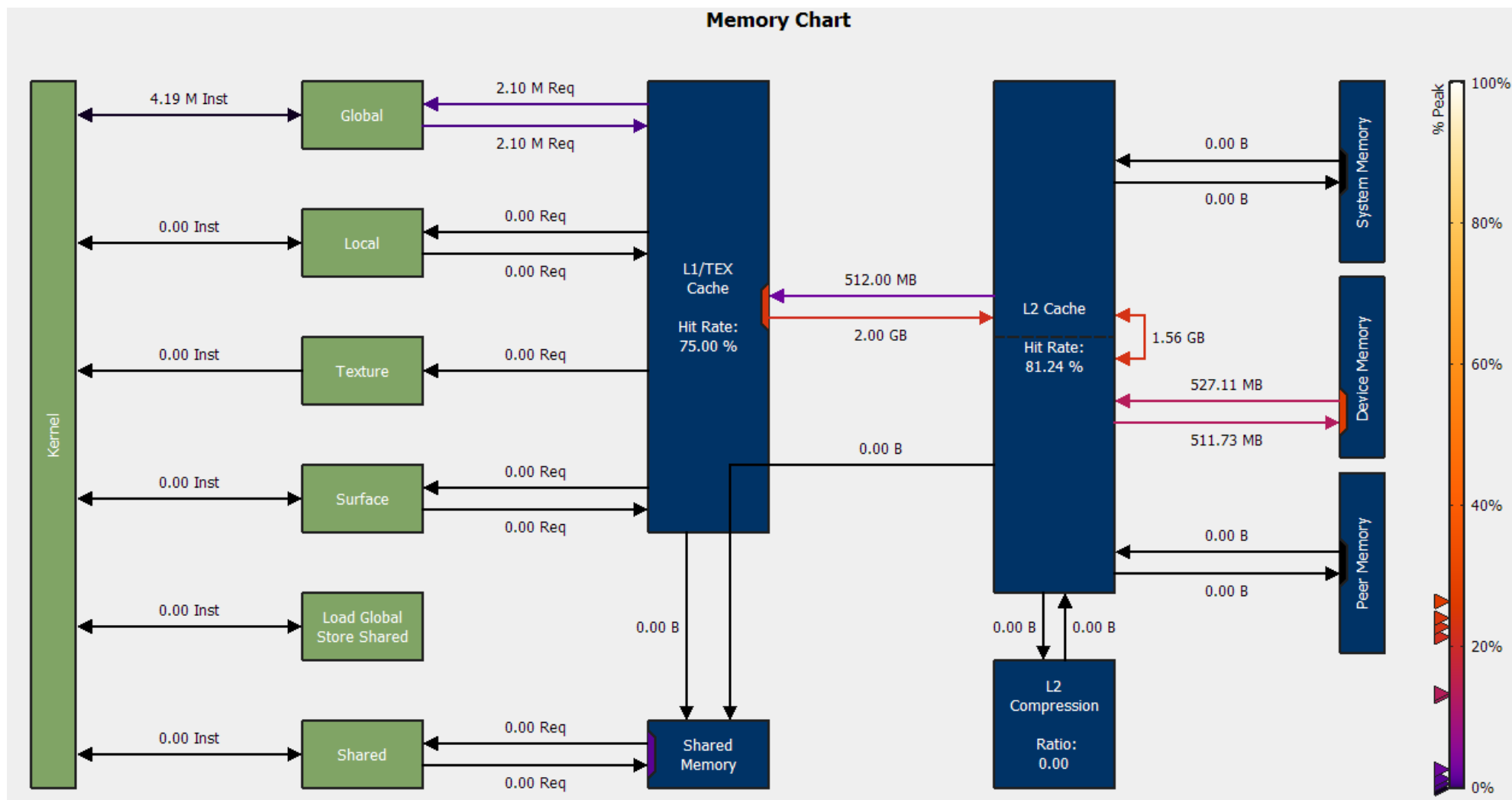
USING NSIGHT COMPUTE (DEMO)



„Uncoalesced global access [...] 4.00x“

GLOBAL, LOCAL, L1, L2?

Understanding the memory hierarchy



MEMORY TRANSACTIONS AND COALESCING

Access to global memory triggers transactions ([Device Mem Access](#))

Memory access granularity = **32 bytes = 1 sector**

Cache line = **128 bytes = 4 consecutive sectors**

Example: 4 byte per thread $\rightarrow 4B * 32$ threads (1 warp) = 128B

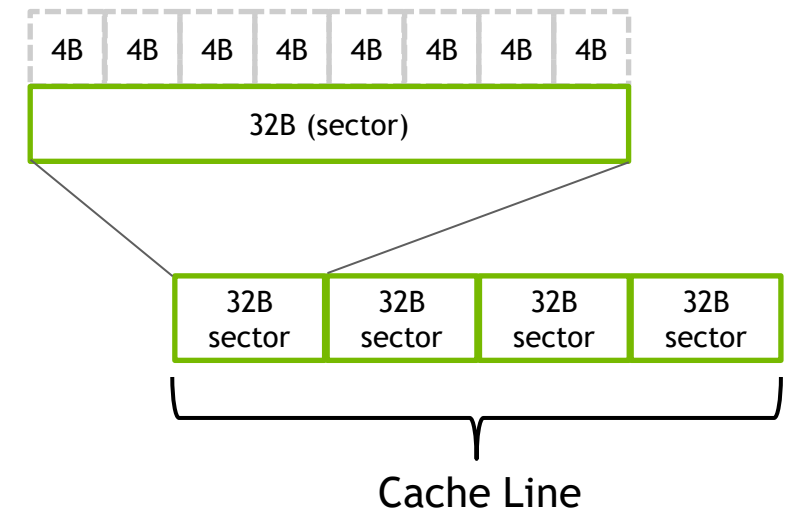
L1-cached access:

128-byte transactions

L2-cached access:

32-byte transactions

Used for global memory



The full picture:
[S32089: Understanding and Optimizing Memory-Bound Kernels with Nsight Compute](#)

MEMORY TRANSACTIONS AND COALESCING

Coalescing details

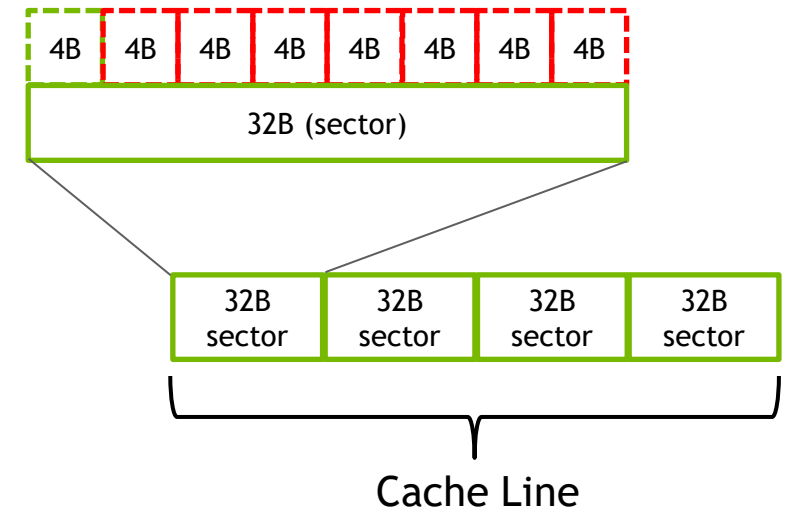
Coalescing: Adjacent threads can share transactions

Transactions must be “Naturally Aligned”: First address % size == 0

All bytes in a transaction are transferred. Use them!

For example, if a 32-byte memory transaction is generated for each thread's 4-byte access, throughput is divided by 8.

$$\text{degree of coalescing} = \frac{\text{\#bytes requested}}{\text{\#bytes transferred}}$$



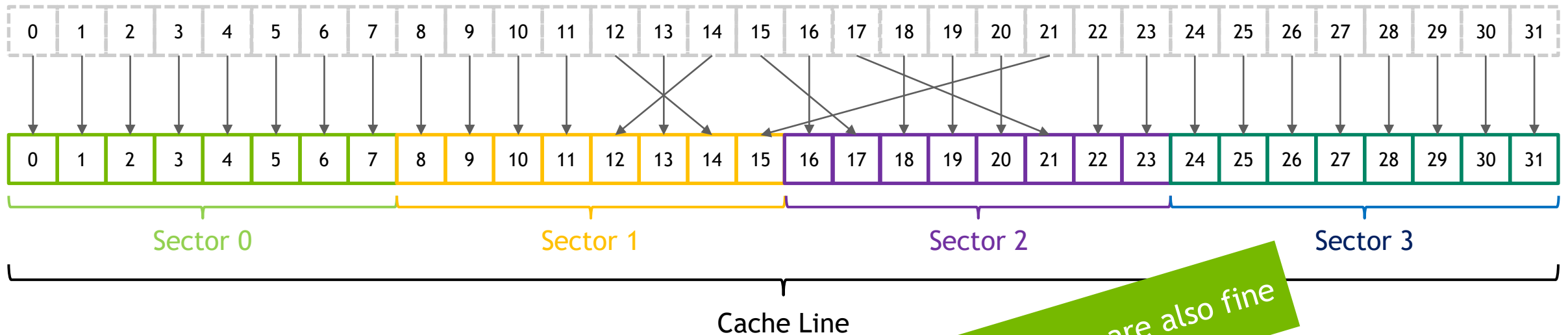
ACCESSING GLOBAL MEMORY

optimal access pattern (4byte words) - fully coalesced

```
int x_val = x[threadIdx.x];
```

All addresses fall within 4 sectors

Bus utilization: 100%



Permutations are also fine

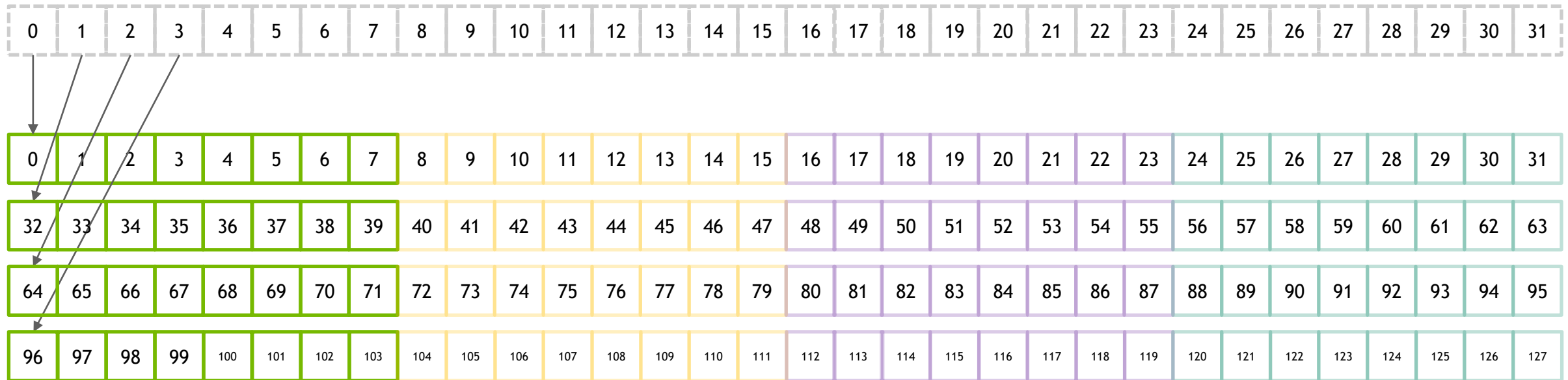
ACCESSING GLOBAL MEMORY

worst case access pattern (4byte words) - fully uncoalesced

```
// stride 32
int x_val = x[32*threadIdx.x];
// „random“ (pointer chasing, lists, tree, ...)
int x_val = x[lookup[threadIdx.x]];
```

All addresses fall in 32 different sectors

Bus utilization: 12.5%



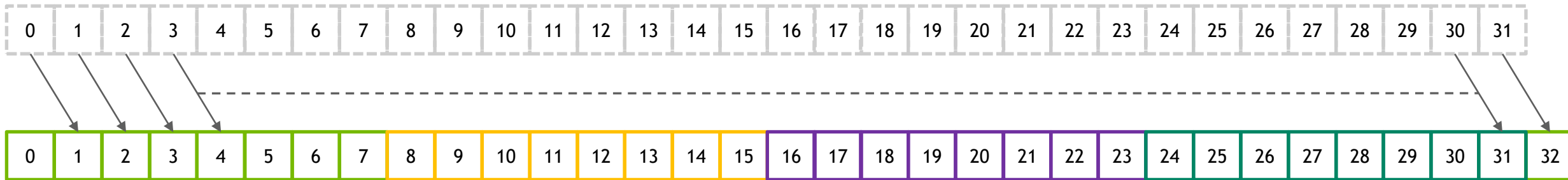
ACCESSING GLOBAL MEMORY

shifted access

```
int x_val = x[threadIdx.x+1];
```

All addresses fall within 5 sectors

Bus utilization: 80% = 128B/160B



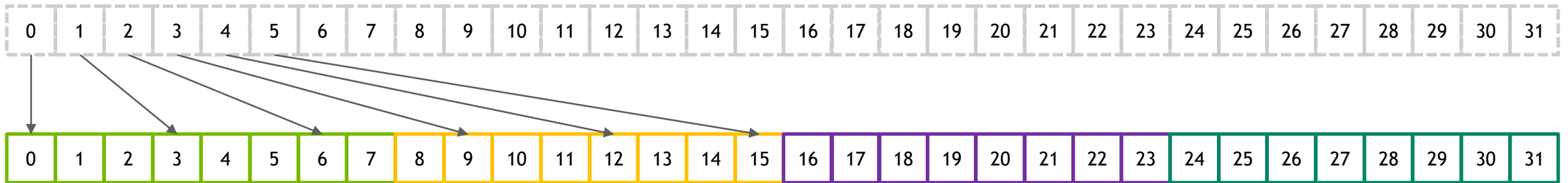
ACCESSING GLOBAL MEMORY

common access pattern: stride 3

```
int x_val = x[3*threadIdx.x];
```

All addresses fall within 12 sectors (4 byte words)

Bus utilization: 33%



```
struct {float x,y,z;} a; ... a[tid].x
```

→ use structure-of-arrays (SoA): `a.x[tid]`

```
float a[M][N]; ... a[tid][42]
```

→ multi-dimensional arrays: pay attention to coalescing (row-major, column-major?)

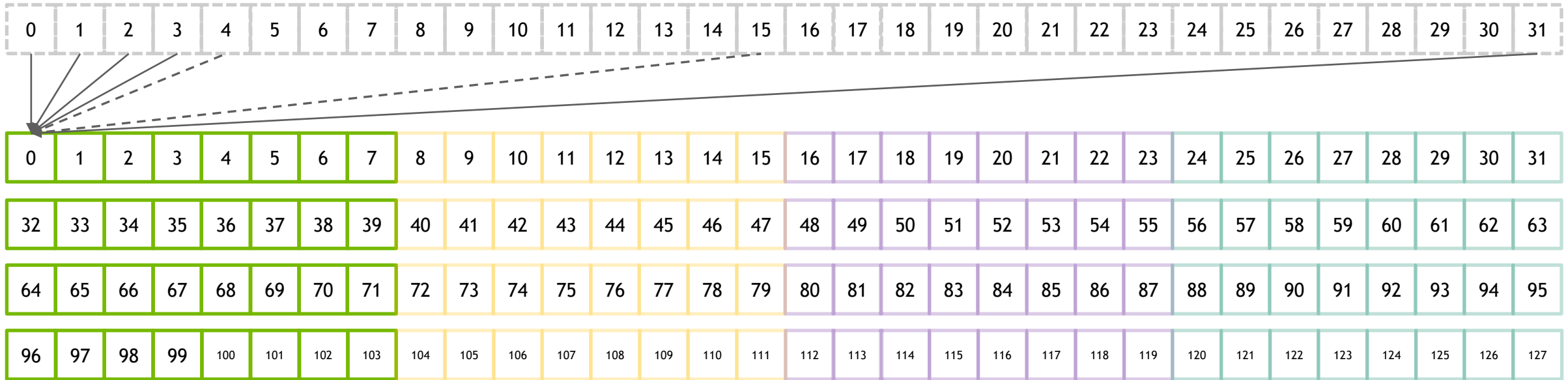
ACCESSING GLOBAL MEMORY

another “worst case” access pattern?

```
// same for all threads - e.g. loop index  
int x_val = x[i];
```

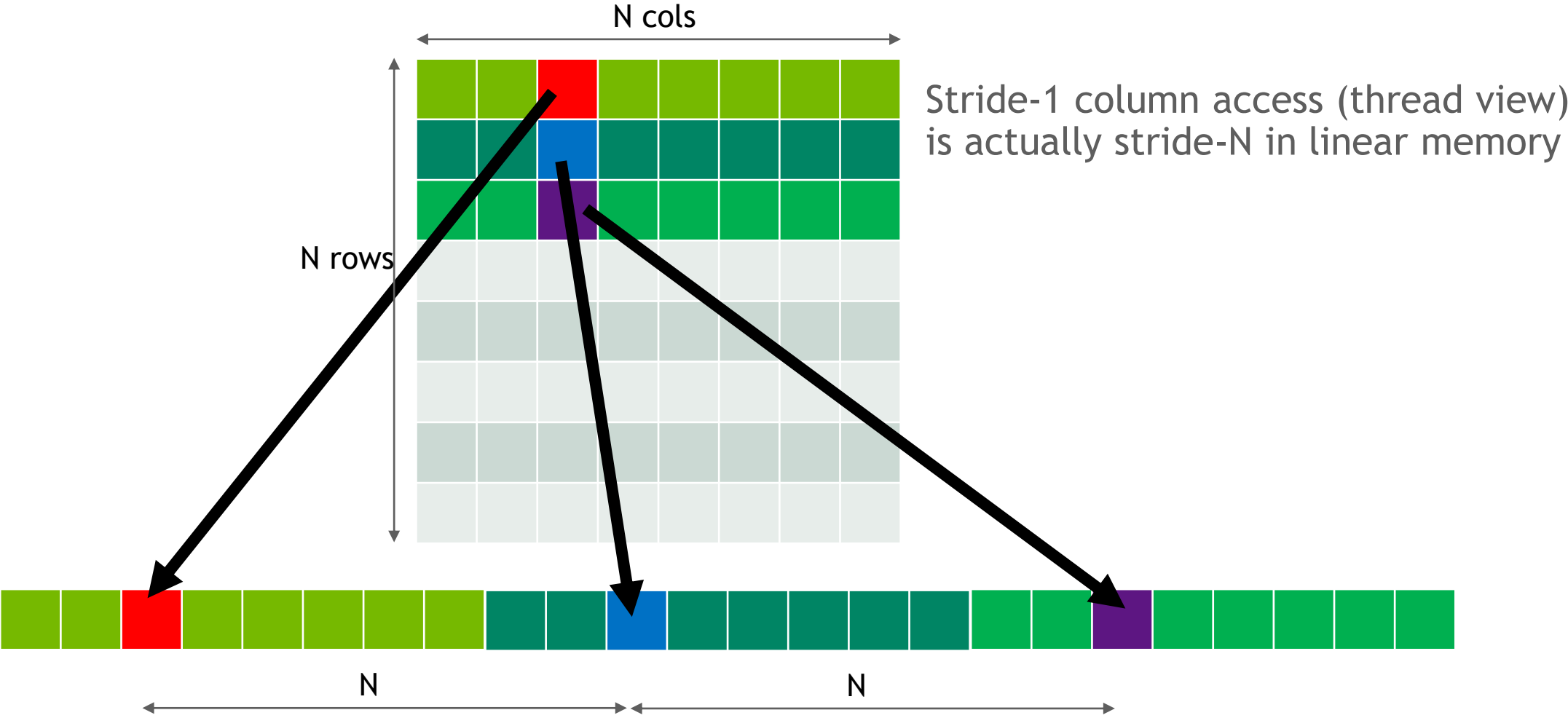
Single address, single sector

Bus utilization: 12.5%



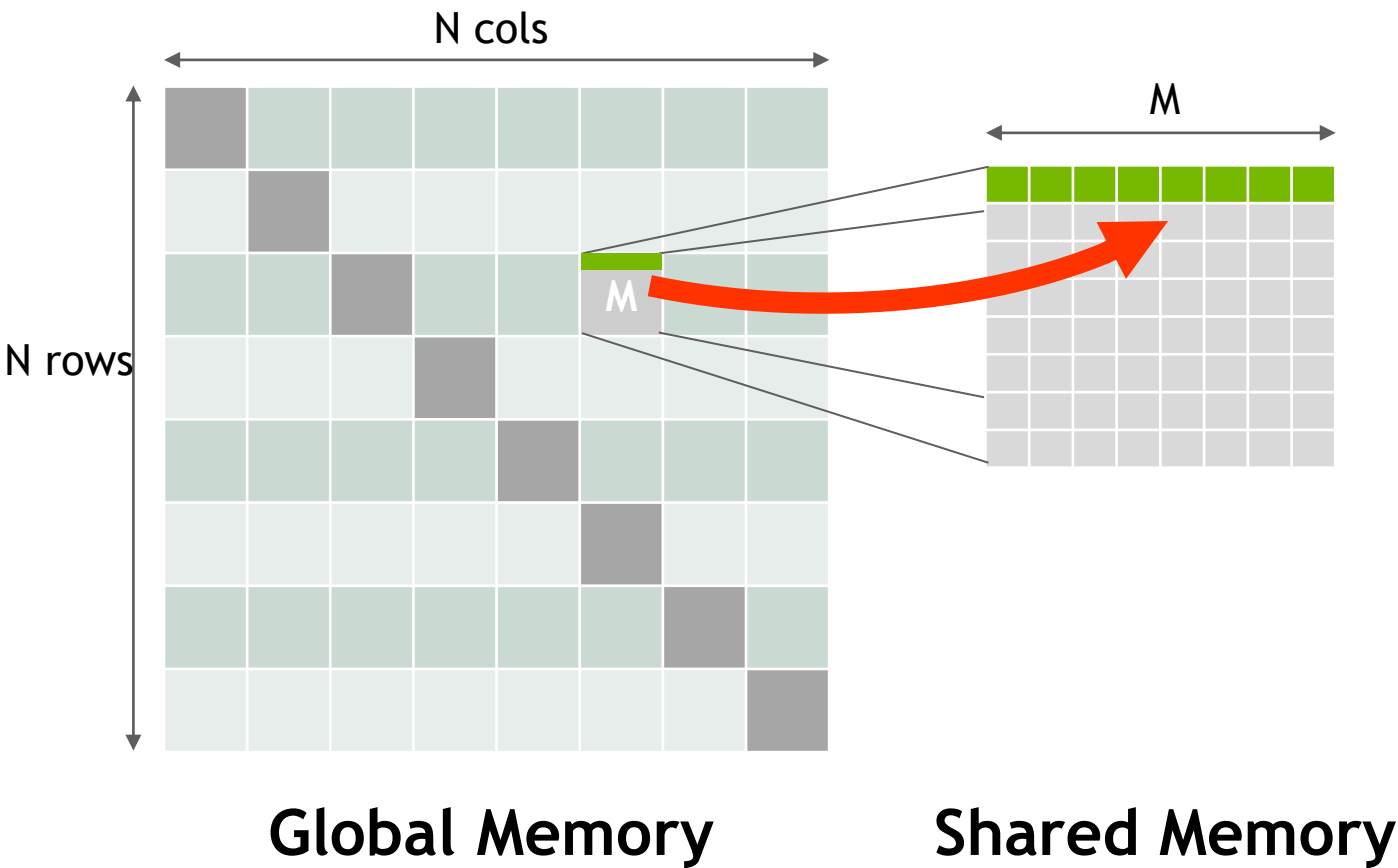
MATRIX TRANSPOSE

Access pattern



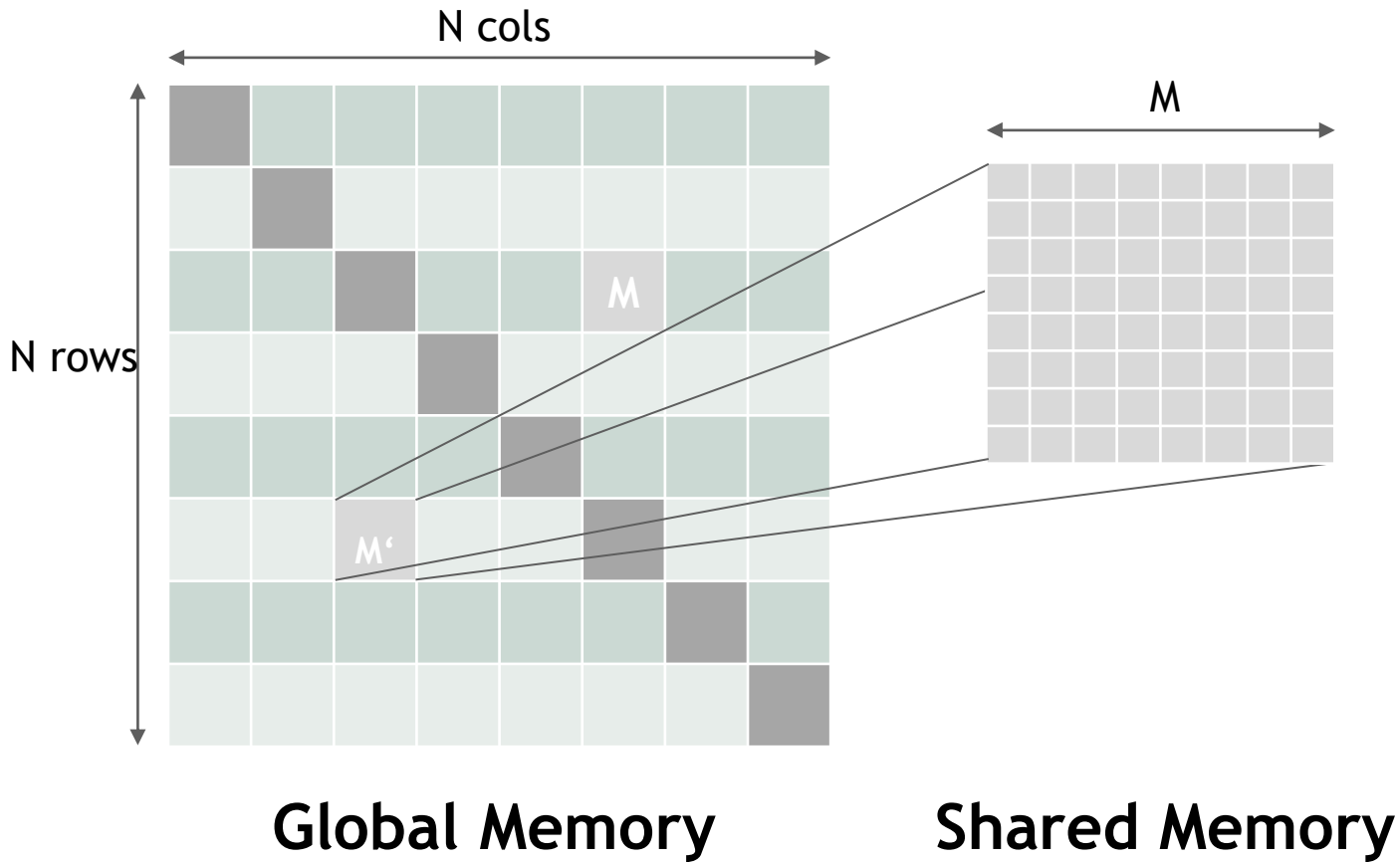
MATRIX TRANSPOSE

Using shared memory



MATRIX TRANSPOSE

Using shared memory

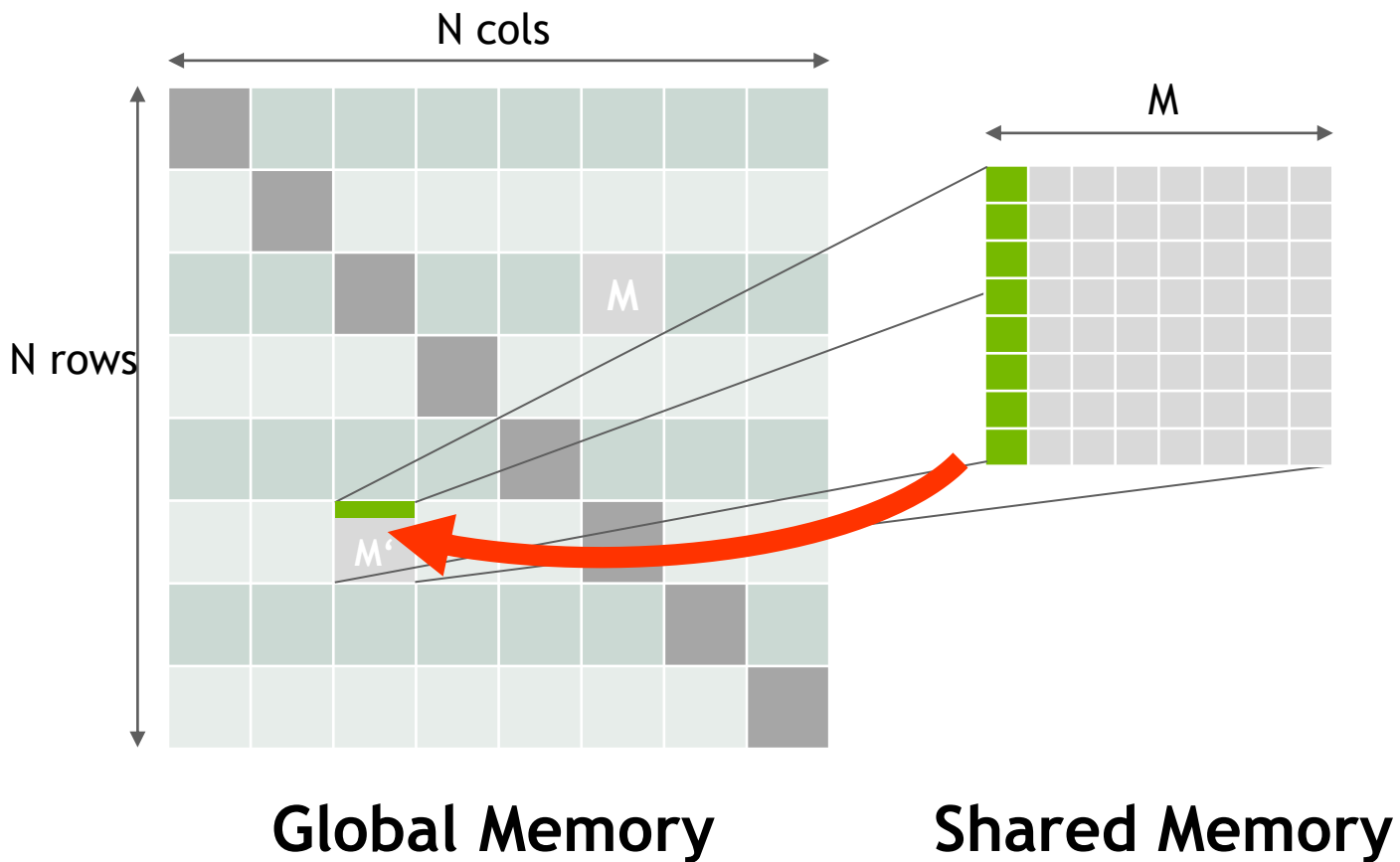


Block row is loaded, fully **coalesced** read

Indexing: Location of block is reflected on the diagonal

MATRIX TRANSPOSE

Using shared memory



Block row is loaded, fully **coalesced** read

Indexing: Location of block is reflected on the diagonal

Block *column* of shared is written to *row* of matrix

Transposes the block $M \rightarrow M'$

Coalesced write

ANALYSIS WITH NSIGHT COMPUTE

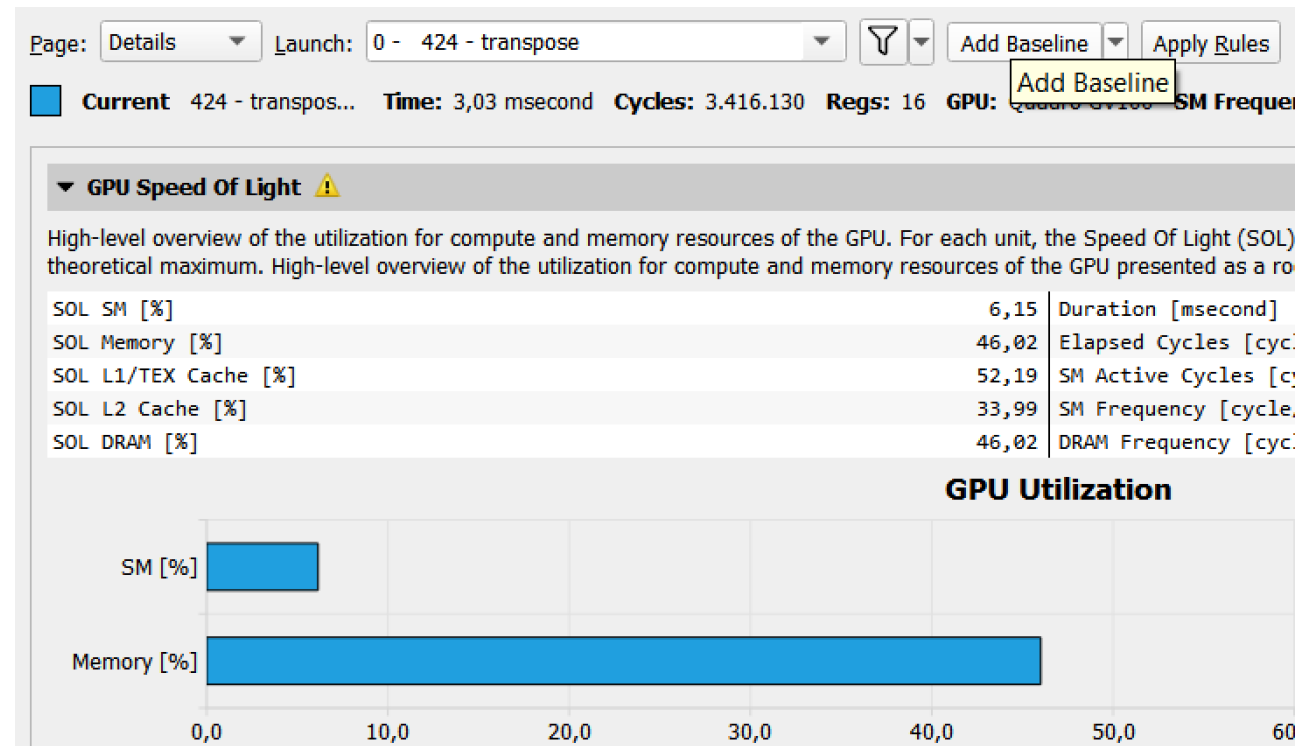
Iterating and comparing

Add baseline for comparison

Check the „Breakdown“ tables and how they change

Look at the other sections, warp state statistics

Tooltips on mouseover over metrics/names/...

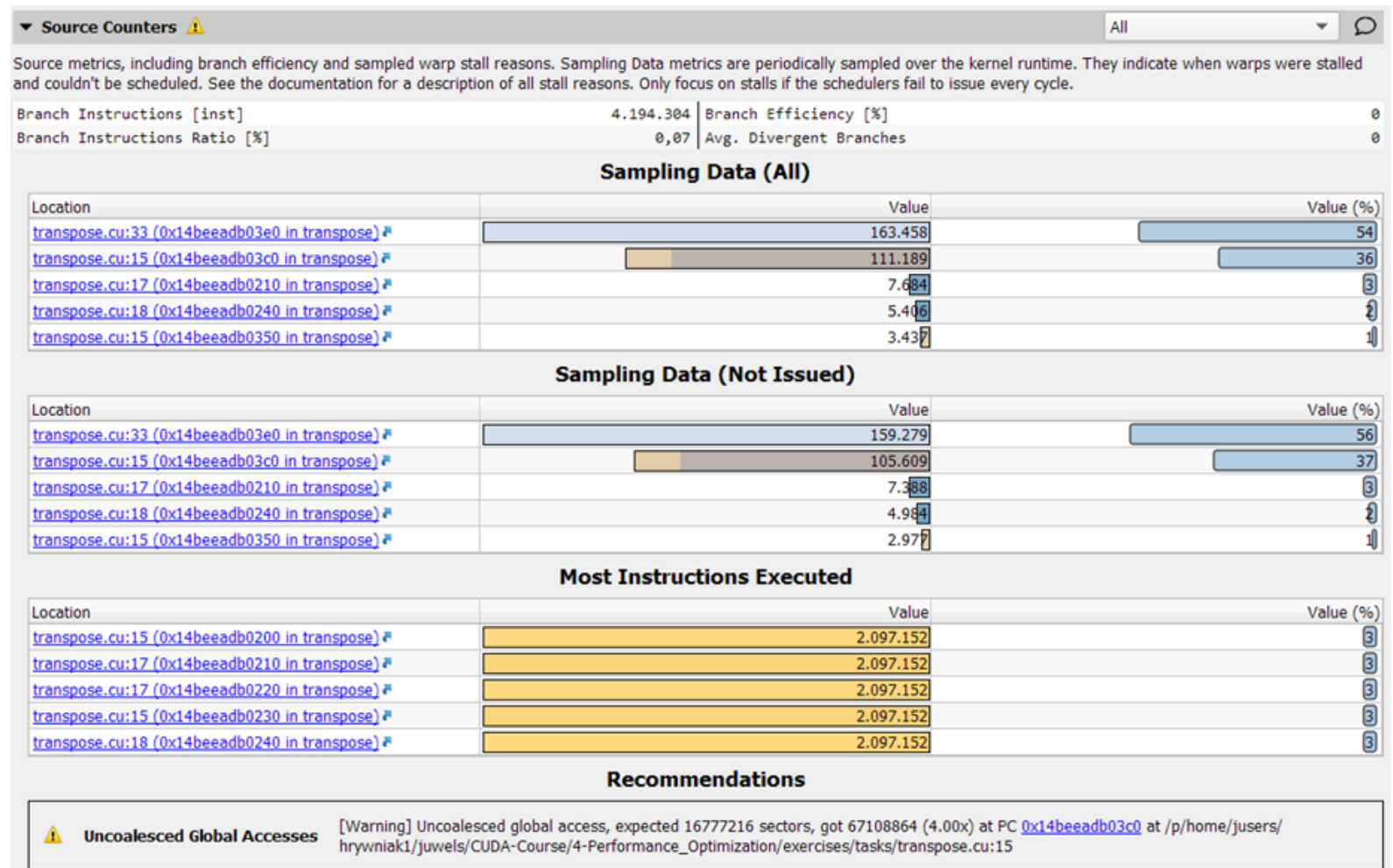


ANALYSIS WITH NSIGHT COMPUTE

Iterating and comparing

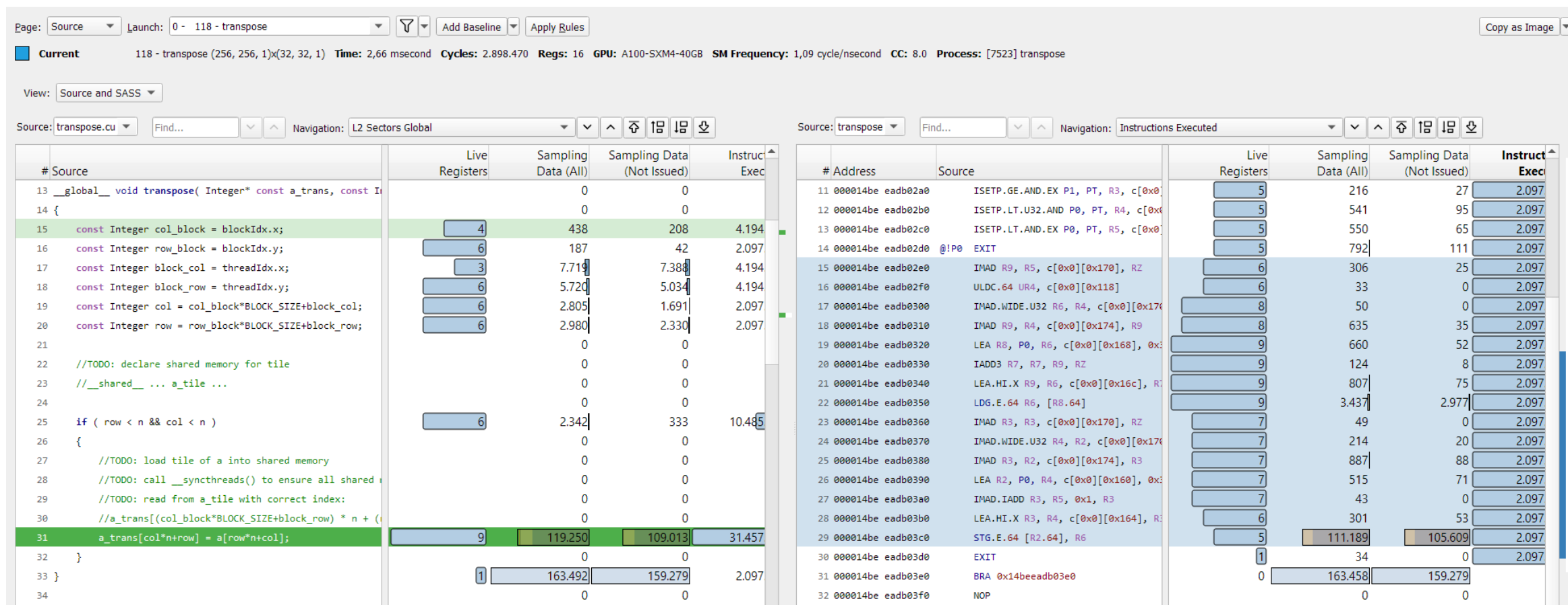
Check the Source Counters section (also on CLI)

Links will take you to Source/SASS view



ANALYSIS WITH NSIGHT COMPUTE

Iterating and comparing



TASK: COALESCED MATRIX TRANSPOSE

Location of code: *4-Performance_Optimization/exercises/tasks*

Steps (see also `Instructions.ipynb`):

Follow TODOs in *transpose.cu*, use shared memory to coalesce writes to *a_trans*

Profile with Nsight Compute and observe your changes

ROOFLINE ANALYSIS

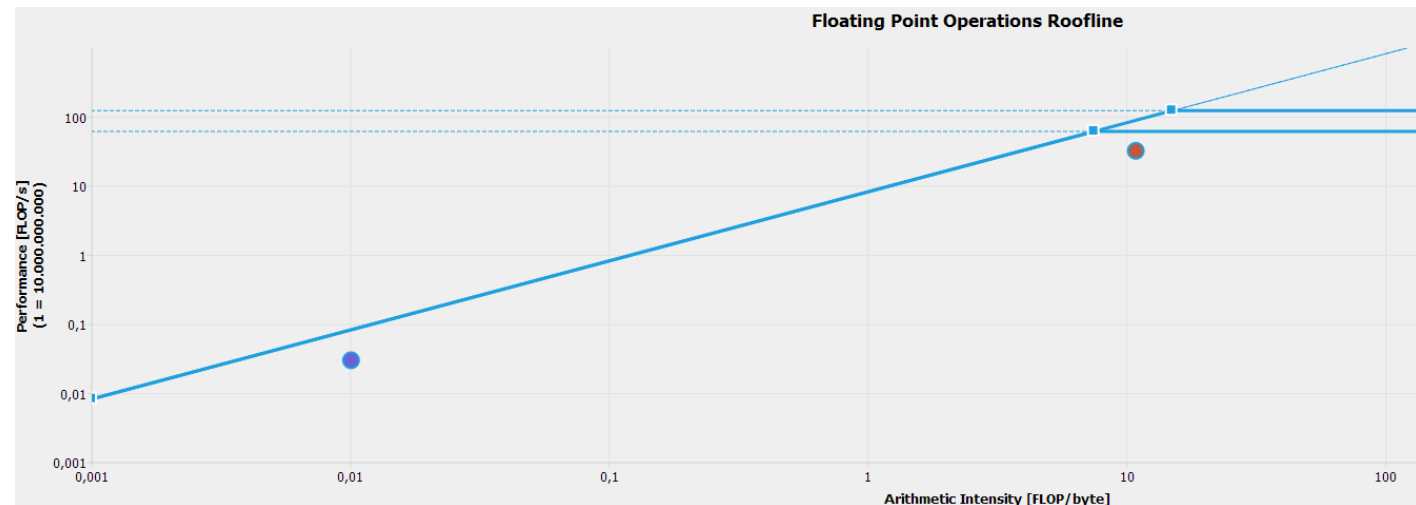
How well is the hardware utilized?

Transpose does zero floating point computations... more interesting example (here: on V100)

Counting flops and transferred bytes $\rightarrow$ AI, x-axis

Measuring achieved performance $\rightarrow$ FLOP/s, y-axis

Rooflines from device peak bandwidth / compute



GTC session:

[S32062: Performance Tuning CUDA Applications with the Roofline Model](#)

Roofline Hackathon:

<https://www.youtube.com/watch?v=ZXZ2SrM3pmE&t=2382s>

BRANCH DIVERGENCE

Recap: Warp execution

GPUs use the Single Instruction Multiple Threads (SIMT) execution

functionally transparent to the programmer

but has performance implications

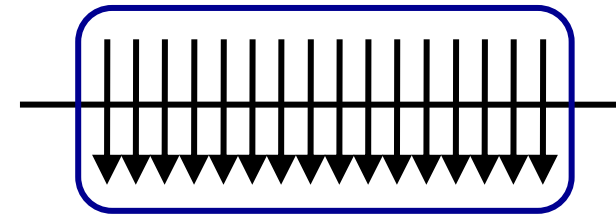
warp

group of synchronously* executing threads

(*since Volta: Independent Thread Scheduling)

neighbor threads (mostly x dimension)

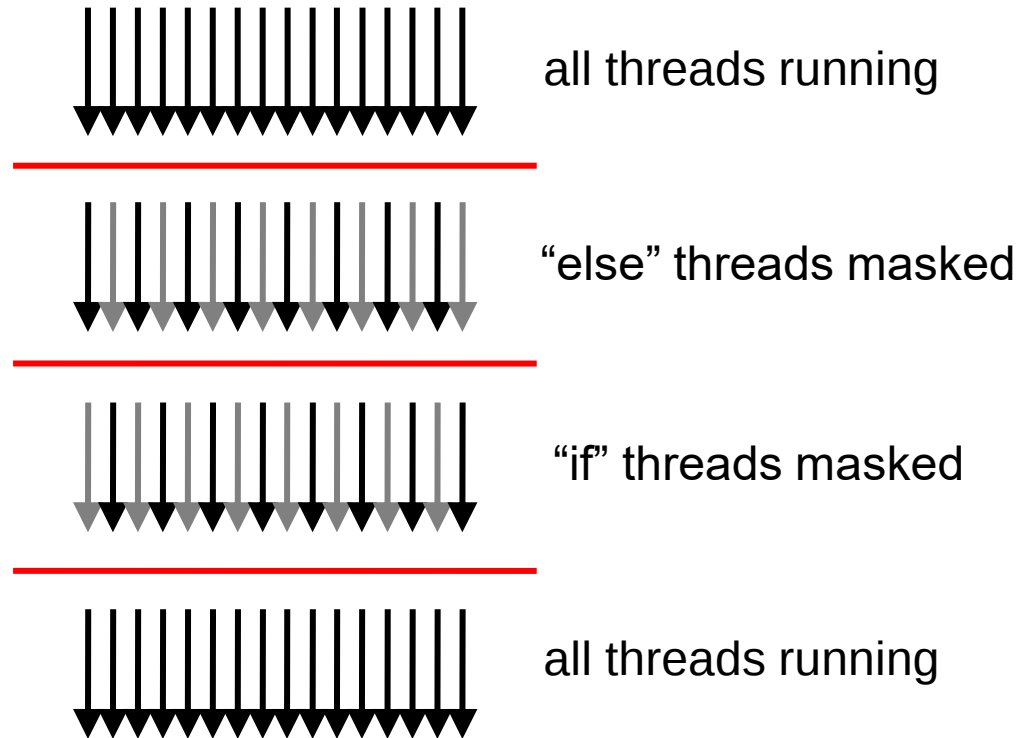
basic unit of scheduling



BRANCH DIVERGENCE

Within Warp

```
// ...  
if (condition) {  
    // do smth  
    // ...  
} else {  
    // do smth else  
    // ...  
}  
// ...
```



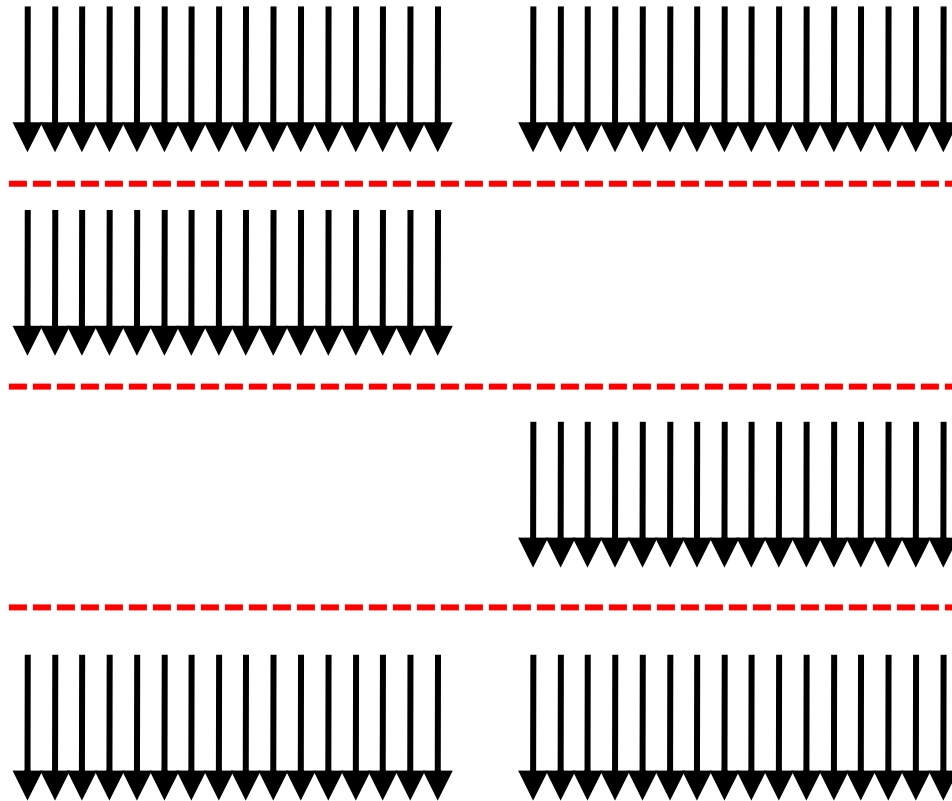
divergence *within* warp → performance penalty

```
if(threadIdx.x % 2 == 0) ...
```

BRANCH DIVERGENCE

Between Warps

```
// ...  
if (condition) {  
    // do smth  
    // ...  
} else {  
    // do smth else  
    // ...  
}  
// ...
```



divergence *between* warps → no penalty

```
if(blockIdx.x % 2 == 0) ...
```

CONCLUSION

To achieve coalesced global memory access:

- Try to use shared memory

- Look for different way of storage or better algorithm

Avoid divergent branches

Use the tools!