



NVIDIA Developer Tools

Kristopher Keipert, Solutions Architect

ATPESC 2015



kkeipert@nvidia.com

Programming the NVIDIA Platform

CPU, GPU, and Network

ACCELERATED STANDARD LANGUAGES

ISO C++, ISO Fortran

```
std::transform(par, x, x+n, y, y,  
    [=](float x, float y){ return y +  
a*x; }  
);
```

```
do concurrent (i = 1:n)  
    y(i) = y(i) + a*x(i)  
enddo
```

```
import cunumeric as np  
...  
def saxpy(a, x, y):  
    y[:] += a*x
```

INCREMENTAL PORTABLE OPTIMIZATION

OpenACC, OpenMP

```
#pragma acc data copy(x,y) {  
...  
std::transform(par, x, x+n, y, y,  
    [=](float x, float y){  
        return y + a*x;  
    });  
...  
}  
  
#pragma omp target data map(x,y) {  
...  
std::transform(par, x, x+n, y, y,  
    [=](float x, float y){  
        return y + a*x;  
    });  
...  
}
```

PLATFORM SPECIALIZATION

CUDA

```
__global__  
void saxpy(int n, float a,  
    float *x, float *y) {  
    int i = blockIdx.x*blockDim.x +  
        threadIdx.x;  
    if (i < n) y[i] += a*x[i];  
}  
  
int main(void) {  
    ...  
    cudaMemcpy(d_x, x, ...);  
    cudaMemcpy(d_y, y, ...);  
  
    saxpy<<<(N+255)/256,256>>>(...);  
  
    cudaMemcpy(y, d_y, ...);  
}
```

ACCELERATION LIBRARIES

Core

Math

Communication

Data Analytics

AI

Quantum

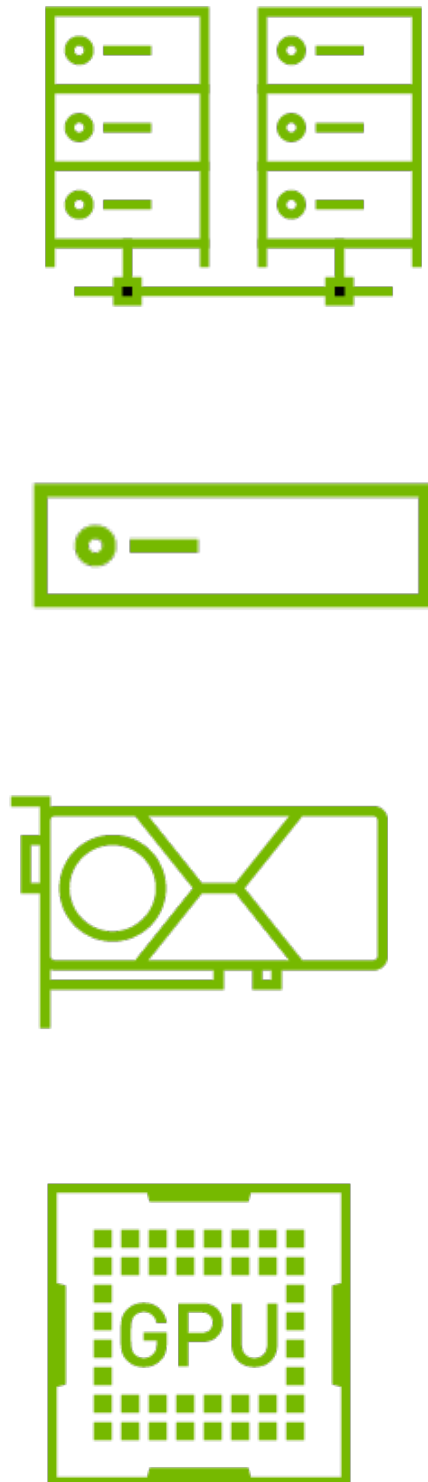
Agenda

- High-level tools ecosystem
- For each tool:
 - Brief description and feature overview
 - What's new in the latest releases
 - Focus on problem solving and application in the system hierarchy
- Summary - Putting it all together

Macro

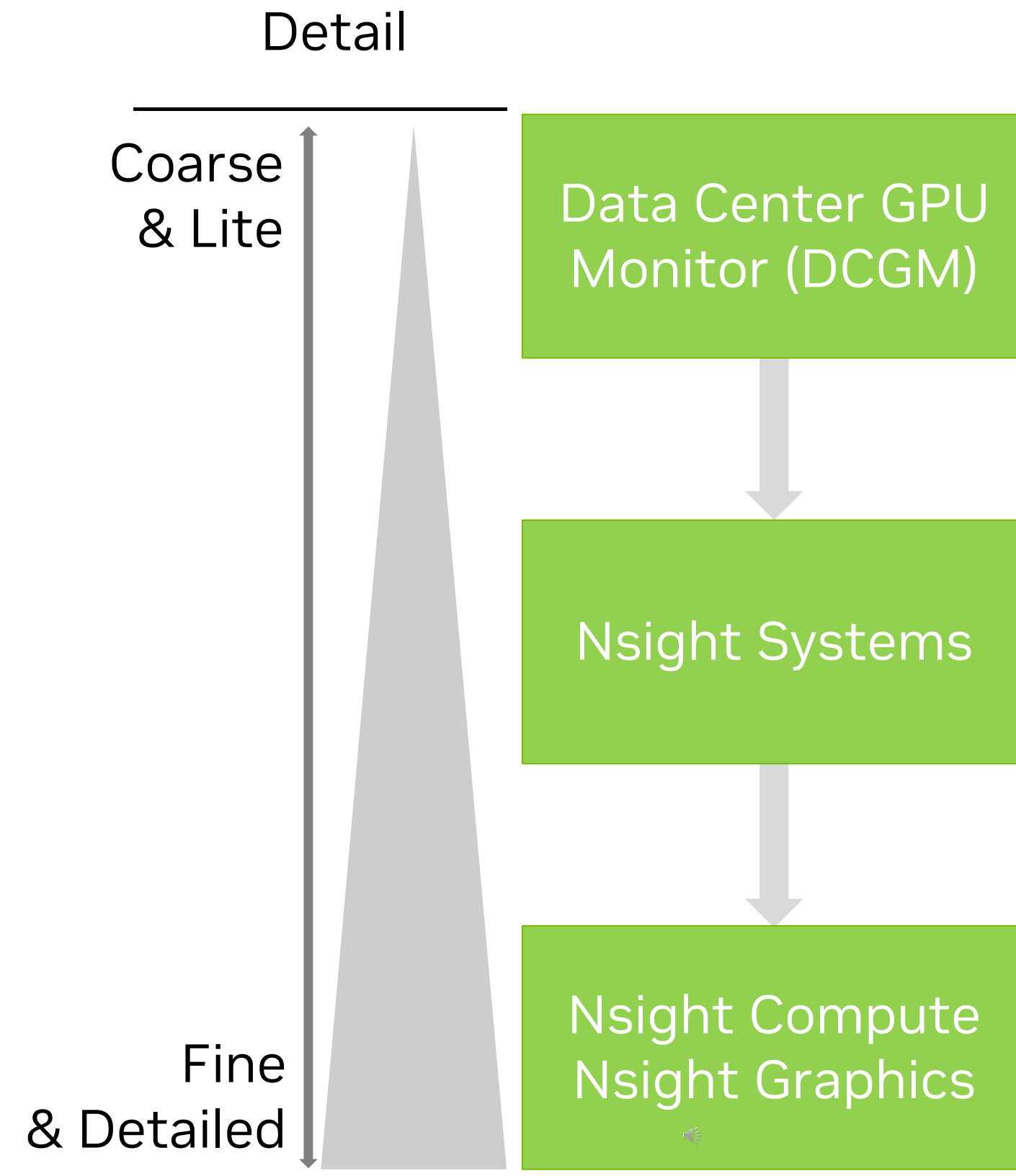


Micro



Monitoring → Profiling

- Different target users, intent, details, design tradeoffs
 - Some overlap but not interchangeable
- Monitors
 - Goal: Coarsely observe quality, utilization, progress, help rate jobs
 - Users: Admins & users of top, netstats, taskman, ...
 - Track issues back to jobs
 - Low rate (minutes or seconds)
 - Sensible to observer checking & reacting, low overhead
 - Less details towards root cause
 - Uptime: Usually 100%
- Profilers
 - Goal: Aid in application optimization
 - Users: Engineers looking to relating back to areas of code
 - Uptime: Usually not 100%



Data Center GPU Monitor (DCGM)

- Very unobtrusive
- Great for fleet health, job health, job rating
- Optimized for low frequency statistics (10hz or less)
- Maps to jobs & processes
 - But not code & algorithms, like profilers!
- Pause/resume to profile
 - CLI
 - `dcgmi profile --pause`
 - `dcgmi profile --resume`
 - API
 - `dcgmProfPause`
 - `dcgmProfResume`



Nsight Profiler Family

Performance Analysis Workflow

Nsight Systems

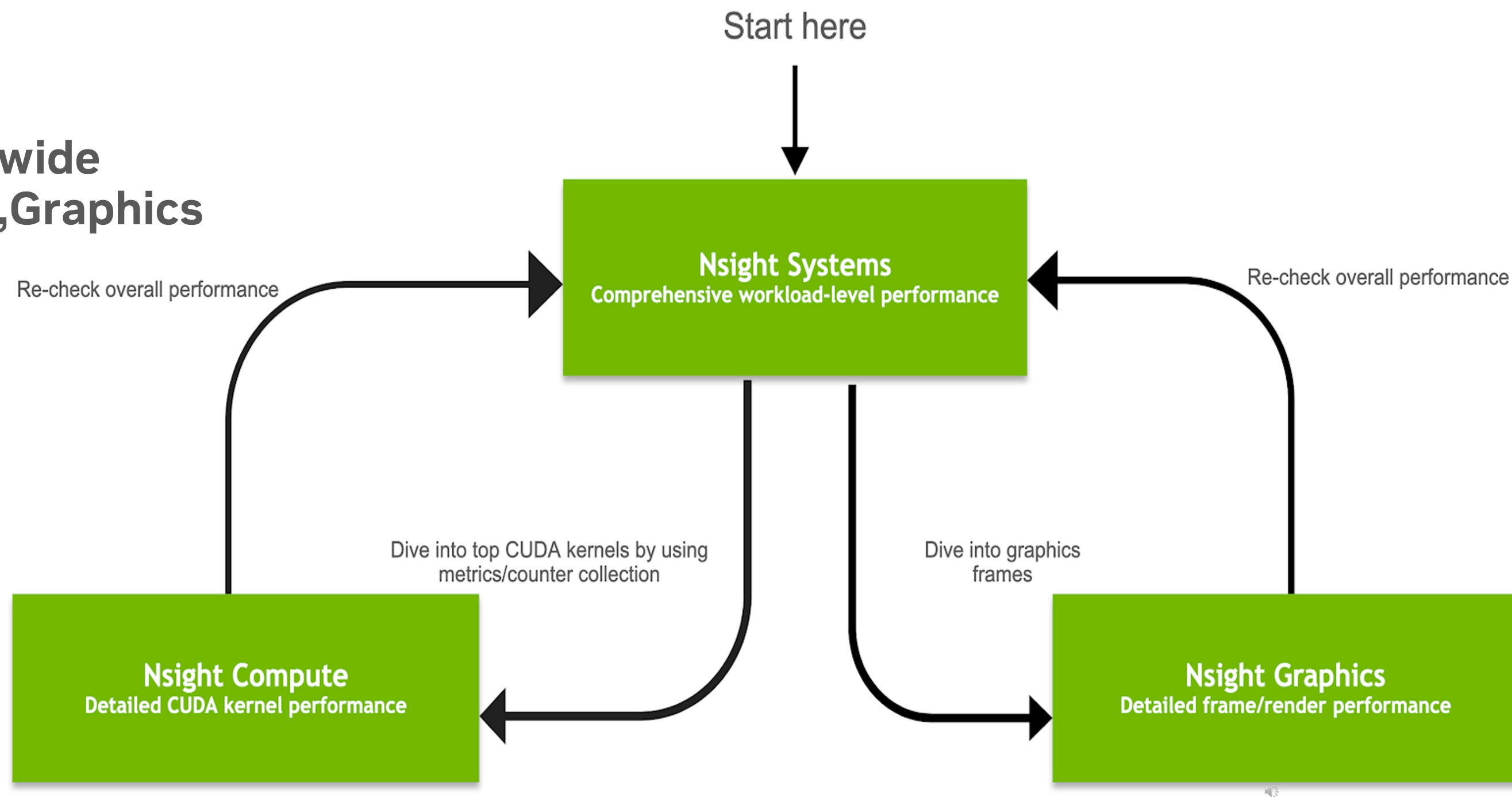
Analyze algorithms system-wide
CPU,GPU,CUDA,Networking,Graphics

Nsight Compute

Debug & analyze
CUDA kernels

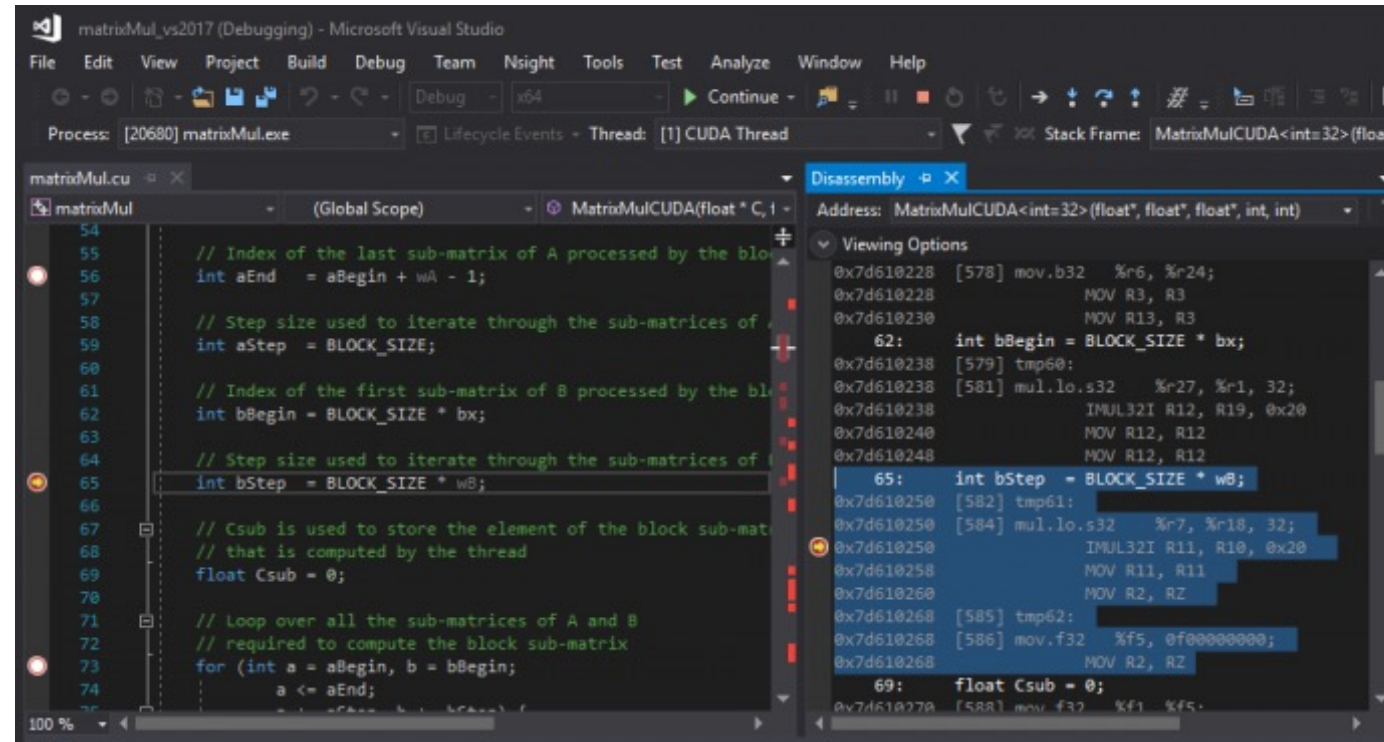
Nsight Graphics

Debug & analyze
Graphics shaders & frames



Developer Tools Ecosystem

Debuggers: cuda-gdb, Nsight Visual Studio Edition
Nsight Visual Studio **Code** Edition



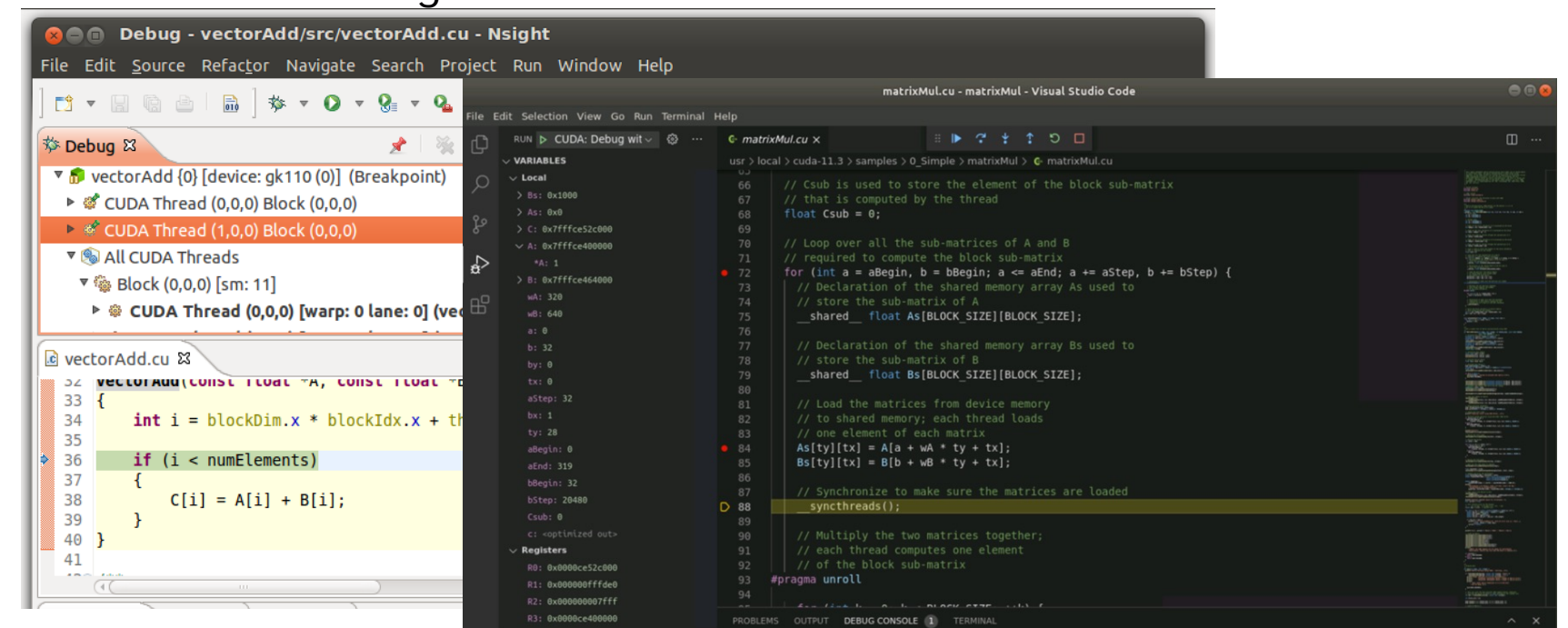
Profilers: Nsight Systems, Nsight Compute, CUPTI, NVIDIA Tools eXtension (NVTX)



Correctness Checker: Compute Sanitizer

```
$ compute-sanitizer --leak-check full memcheck_demo
===== COMPUTE-SANITIZER
Mallocing memory
Running unaligned_kernel
Ran unaligned_kernel: no error
Sync: no error
Running out_of_bounds_kernel
Ran out_of_bounds_kernel: no error
Sync: no error
===== Invalid __global__ write of size 4 bytes
===== at 0x60 in memcheck_demo.cu:6:unaligned_kernel(void)
===== by thread (0,0,0) in block (0,0,0)
===== Address 0x400100001 is misaligned
```

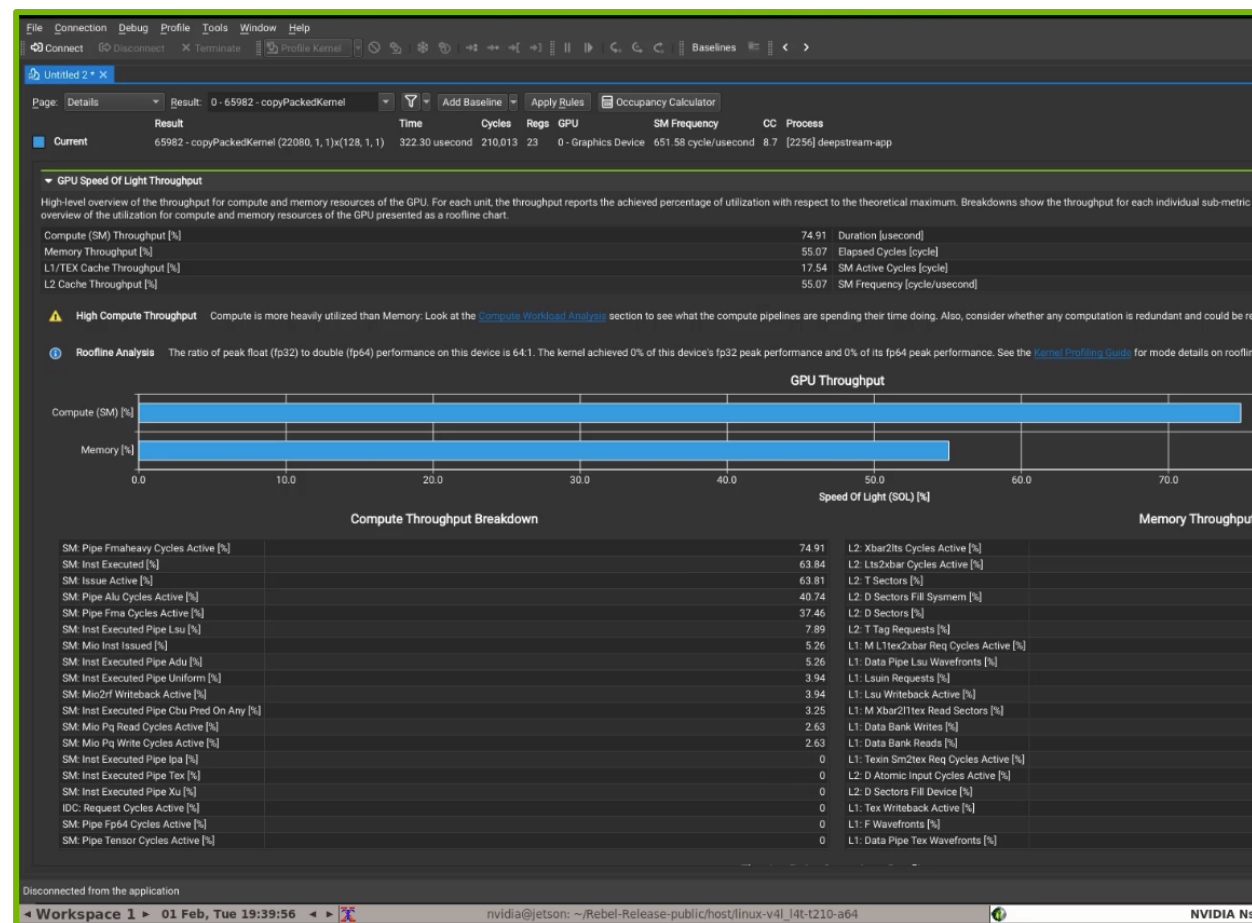
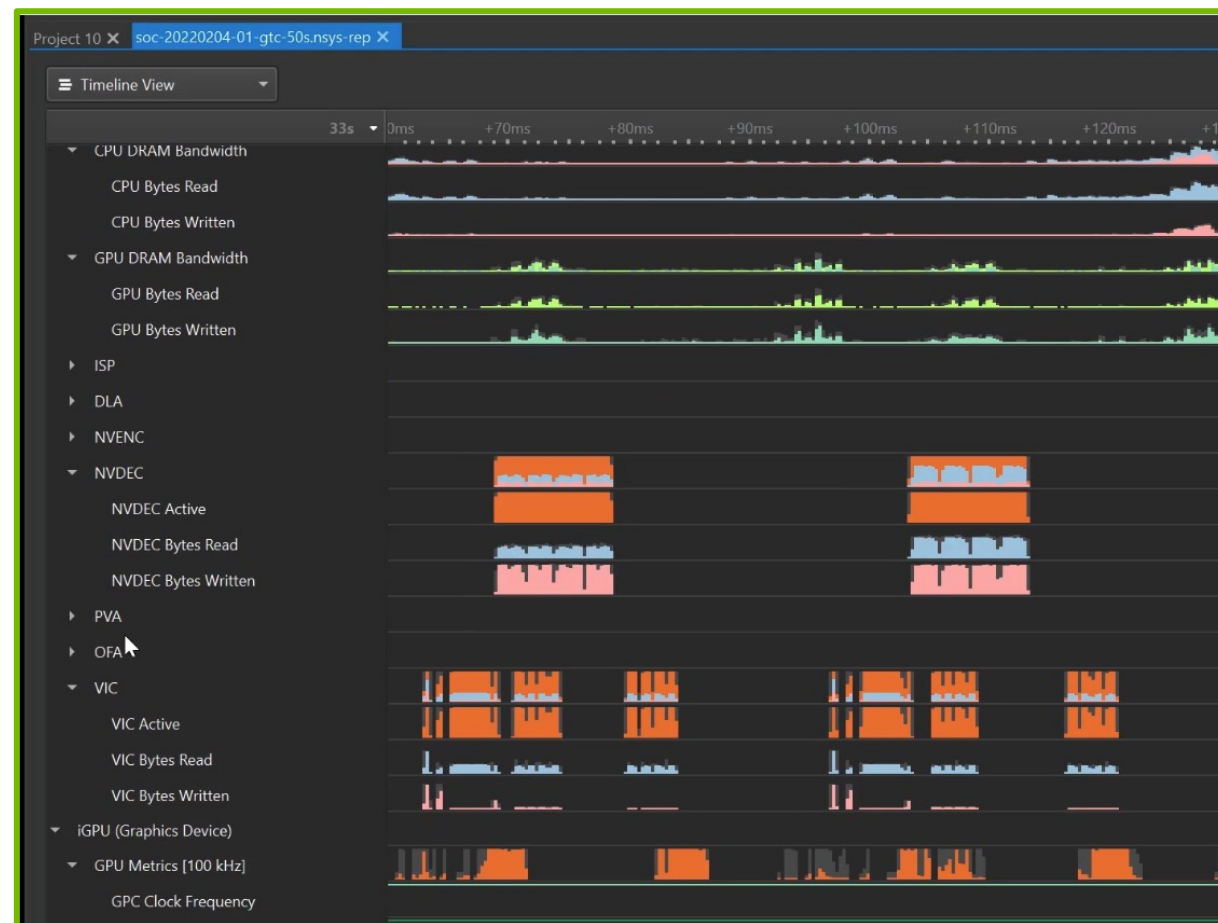
IDE integrations: Nsight Eclipse Edition
Nsight Visual Studio Edition
Nsight Visual Studio **Code** Edition



GUI Support for Grace and other Arm Platforms

Nsight Systems - Nsight Compute - Nsight Visual Studio Code Edition

- GUIs run natively on
 - NVIDIA® Jetson AGX Orin™ SoC
 - Arm Server Platforms
- Use new native GUI or existing remote collection capabilities



The screenshot displays the Visual Studio Code editor with the 'matrixMul.cu' file open. The code is written in C++ and includes CUDA-specific headers and functions. The code is organized into sections: 'Local' variables, 'WATCH' section, 'CALL STACK' section, and 'BREAKPOINTS' section. The 'WATCH' section shows the value of 'a' as '<not available>' and 'b' as '<not available>'. The 'CALL STACK' section shows the function 'main' at line 275. The 'BREAKPOINTS' section shows a breakpoint at line 275, 'main'.

Updated Websites/Docs/Icons

NVIDIA Developer Tools

NVIDIA Nsight™ tools are a powerful set of libraries, SDKs, and developer tools spanning across desktop and mobile targets that enable developers to build, debug, profile, and develop software that utilizes the latest accelerated computing hardware.

	IDE / DEBUGGING						PROFILERS / DESIGN				LIBRARIES / API / SDK			
	Nsight Visual Studio Edition	Nsight Eclipse Edition	Nsight Visual Studio Code Edition	CUDA-GDB	Compute Sanitizer	Nsight Graphics	Nsight Compute	Nsight Systems	Nsight Deep Learning Designer	Nsight Graphics	CUPTI	NVTX	Nsight Aftermath SDK	Nsight Perf SDK
CUDA	●	●	●	●	●		●	●	●		●	●		
GRAPHICS						●		●	●	●		●	●	●
OPTIX	●	●	●	●	●		●	●			●	●		

What Segment Are You Working on?

[View All](#) [CUDA/Compute](#) [Graphics](#) [OptiX](#) [Deep Learning](#)

Developer Activity

CODE DEVELOPMENT

DEBUGGING / CORRECTNESS

PROFILING

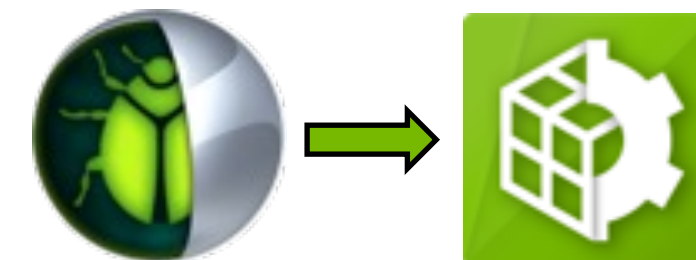
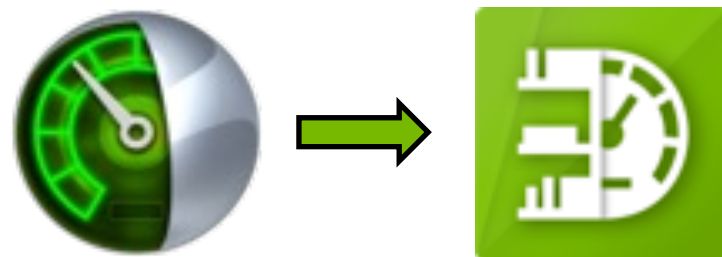
PLATFORM ANALYSIS

KERNEL ANALYSIS

APIs

Nsight Compute

Nsight Compute is an interactive kernel profiler for CUDA applications. It provides detailed performance metrics and API



<https://developer.nvidia.com/tools-overview>

The background features a complex pattern of glowing green lines of varying thicknesses and colors (ranging from light green to dark green/black) against a solid black background. These lines are mostly oriented diagonally, creating a sense of motion and depth. On the far left, there is a solid, bright green vertical rectangular bar.

Compute Debuggers/IDEs

Compute Debuggers

Debug GPU Kernels Running on Device



- CUDA GDB
 - CPU + GPU CUDA kernel debugger
 - Supports stepping, breakpoints, in-line functions, variable inspection etc...
- Nsight Visual Studio Edition
 - IDE integration for Visual Studio
 - Build and Debug CPU+GPU code from Visual Studio
- Nsight Visual Studio Code Edition
 - New IDE integration for VS Code
 - Build and Debug CPU+GPU code from Visual Studio Code
 - Remotely target Linux targets from Windows or Linux

The screenshot displays the Visual Studio Code interface for debugging a CUDA kernel named `matrixMul.cu`. The left sidebar contains several panels: **VARIABLES** (showing local variables like `Bs`, `As`, `C`, `A`, `a`, `b`, `tx`, `ty`, `aBegin`, `aEnd`, `bBegin`, `bStep`, `Csub`, and registers `R0`, `R1`, `R2`, `R3`), **WATCH** (showing `As: [32]` and `threadIdx: {...}`), and **BREAKPOINTS** (listing breakpoints for `matrixMul.cu` at lines 61, 64, 72, and 84). The central editor shows the C++ code for `matrixMul.cu`, with a breakpoint set at line 84. The **DEBUG CONSOLE** at the bottom displays the output of the debugger, including the command `-exec info cuda warps` and a table of warp information.

Wp	Active	Lanes	Mask	Divergent	Lanes	Mask	Active	Physical	PC	Kernel	BlockIdx	First	Active
Device 0	SM 1												
0		0xffffffff			0x00000000	0x00000000	0x00000000	0x00000000	0x00000000	0 (1,0,0)			
1		0xffffffff			0x00000000	0x00000000	0x00000000	0x00000000	0x00000000	0 (1,0,0)			
2		0xffffffff			0x00000000	0x00000000	0x00000000	0x00000000	0x00000000	0 (1,0,0)			
3		0xffffffff			0x00000000	0x00000000	0x00000000	0x00000000	0x00000000	0 (1,0,0)			
4		0xffffffff			0x00000000	0x00000000	0x00000000	0x00000000	0x00000000	0 (1,0,0)			
5		0xffffffff			0x00000000	0x00000000	0x00000000	0x00000000	0x00000000	0 (1,0,0)			
6		0xffffffff			0x00000000	0x00000000	0x00000000	0x00000000	0x00000000	0 (1,0,0)			
7		0xffffffff			0x00000000	0x00000000	0x00000000	0x00000000	0x00000000	0 (1,0,0)			

- Unified CPU and CUDA Debugging
- CUDA-C/SASS support
- Built on GDB and uses many of the same CLI commands
- Local/Remote connection support
- Backend for IDE debuggers

(cuda-gdb) info cuda threads breakpoint all										
BlockIdx	ThreadIdx	Virtual PC	Dev	SM	Wp	Ln	Filename	Line		
Kernel 0										
(1,0,0)	(0,0,0)	0x00000000000948e58	0	11	0	0	infoCommands.cu	12		
(1,0,0)	(1,0,0)	0x00000000000948e58	0	11	0	1	infoCommands.cu	12		
(1,0,0)	(2,0,0)	0x00000000000948e58	0	11	0	2	infoCommands.cu	12		
(1,0,0)	(3,0,0)	0x00000000000948e58	0	11	0	3	infoCommands.cu	12		
(1,0,0)	(4,0,0)	0x00000000000948e58	0	11	0	4	infoCommands.cu	12		
(1,0,0)	(5,0,0)	0x00000000000948e58	0	11	0	5	infoCommands.cu	12		
(cuda-gdb) info cuda threads breakpoint 2 lane 1										
BlockIdx	ThreadIdx	Virtual PC	Dev	SM	Wp	Ln	Filename	Line		
Kernel 0										
(1,0,0)	(1,0,0)	0x00000000000948e58	0	11	0	1	infoCommands.cu	12		

```
117  /**
118   * Run a simple test of matrix multiplication using CUDA
119   */
120  int MatrixMultiply(int argc, char **argv, int block_size, const dim3 &dimsA,
121                    const dim3 &dimsB) {
122      // Allocate host memory for matrices A and B
123      unsigned int size_A = dimsA.x * dimsA.y;
```

PROBLEMS OUTPUT DEBUG CONSOLE TERMINAL Filter (e.g. text, !exclude)

96 Csub += As[ty][k] * Bs[k][tx];
cuda sm 0 warp 1 lane 0
CUDA focus unchanged.

Thread 1 "matrixMul" hit Breakpoint 2, MatrixMulCUDA<32><<<(20,10,1),(32,32,1)>>> (C=0x7ffffcb2c000, A=0x7ffffcca00000, B=0x7ffffcca64000, wA=320, wB=640) at matrixMul.cu:96
96 Csub += As[ty][k] * Bs[k][tx];

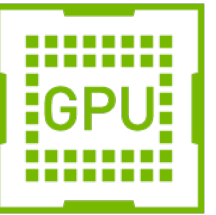
Thread 1 "matrixMul" hit Breakpoint 3, MatrixMulCUDA<32><<<(20,10,1),(32,32,1)>>> (C=0x7ffffcb2c000, A=0x7ffffcca00000, B=0x7ffffcca64000, wA=320, wB=640) at matrixMul.cu:108
108 C[c + wB * ty + tx] = Csub;

>

⊗ ⊕ 0 ⚙ CUDA: Debug with CUDA-GDB (matrixMul) Ln 107, Col 1 Spaces: 2 UTF-8 LF Cuda C++

Compute Sanitizer

Automatically Scan for Bugs and Memory Issues



- Compute Sanitizer checks correctness issues via sub-tools:
 - **Memcheck** – Memory access error and leak detection tool.
 - **Racecheck** – Shared memory data access hazard detection tool.
 - **Initcheck** – Uninitialized device global memory access detection tool.
 - **Synccheck** – Thread synchronization hazard detection tool.

<https://github.com/NVIDIA/compute-sanitizer-samples>

```
$ make run_memcheck
/usr/local/cuda/compute-sanitizer/compute-sanitizer --destroy-on-device-error kernel memcheck_demo
===== COMPUTE-SANITIZER
Mallocing memory
===== Invalid __global__ write of size 4 bytes
===== at 0x70 in unaligned_kernel()
===== by thread (0,0,0) in block (0,0,0)
===== Address 0x7f671ac00001 is misaligned
===== and is inside the nearest allocation at 0x7fb654c00000 of size 4 bytes
===== Saved host backtrace up to driver entry point at kernel launch time
===== Host Frame: [0x2774ec]
===== in /lib/x86_64-linux-gnu/libcuda.so.1
===== Host Frame: __cudart803 [0xfccb]
===== in /home/cuda/github/compute-sanitizer-samples/Memcheck/memcheck_demo
===== Host Frame: cudaLaunchKernel [0x6a578]
===== in /home/cuda/github/compute-sanitizer-samples/Memcheck/memcheck_demo
===== Host Frame: cudaError cudaLaunchKernel<char>(char const*, dim3, dim3, void**, unsigned
===== in /home/cuda/github/compute-sanitizer-samples/Memcheck/memcheck_demo
===== Host Frame: __device_stub__Z16unaligned_kernelv() [0xb22e]
===== in /home/cuda/github/compute-sanitizer-samples/Memcheck/memcheck_demo
===== Host Frame: unaligned_kernel() [0xb28c]
===== in /home/cuda/github/compute-sanitizer-samples/Memcheck/memcheck_demo
===== Host Frame: run_unaligned() [0xaf55]
===== in /home/cuda/github/compute-sanitizer-samples/Memcheck/memcheck_demo
===== Host Frame: main [0xb0e2]
===== in /home/cuda/github/compute-sanitizer-samples/Memcheck/memcheck_demo
===== Host Frame: ../sysdeps/nptl/libc_start_call_main.h:58: __libc_start_call_main [0x2dfd0]
===== in /lib/x86_64-linux-gnu/libc.so.6
```


The background of the slide is a black field filled with numerous thin, curved, and straight lines in shades of green and white. These lines create a sense of motion and depth, resembling a stylized representation of data or a network. The lines are most concentrated in the lower right quadrant, where they form a complex, almost crystalline structure, and become more sparse towards the top left.

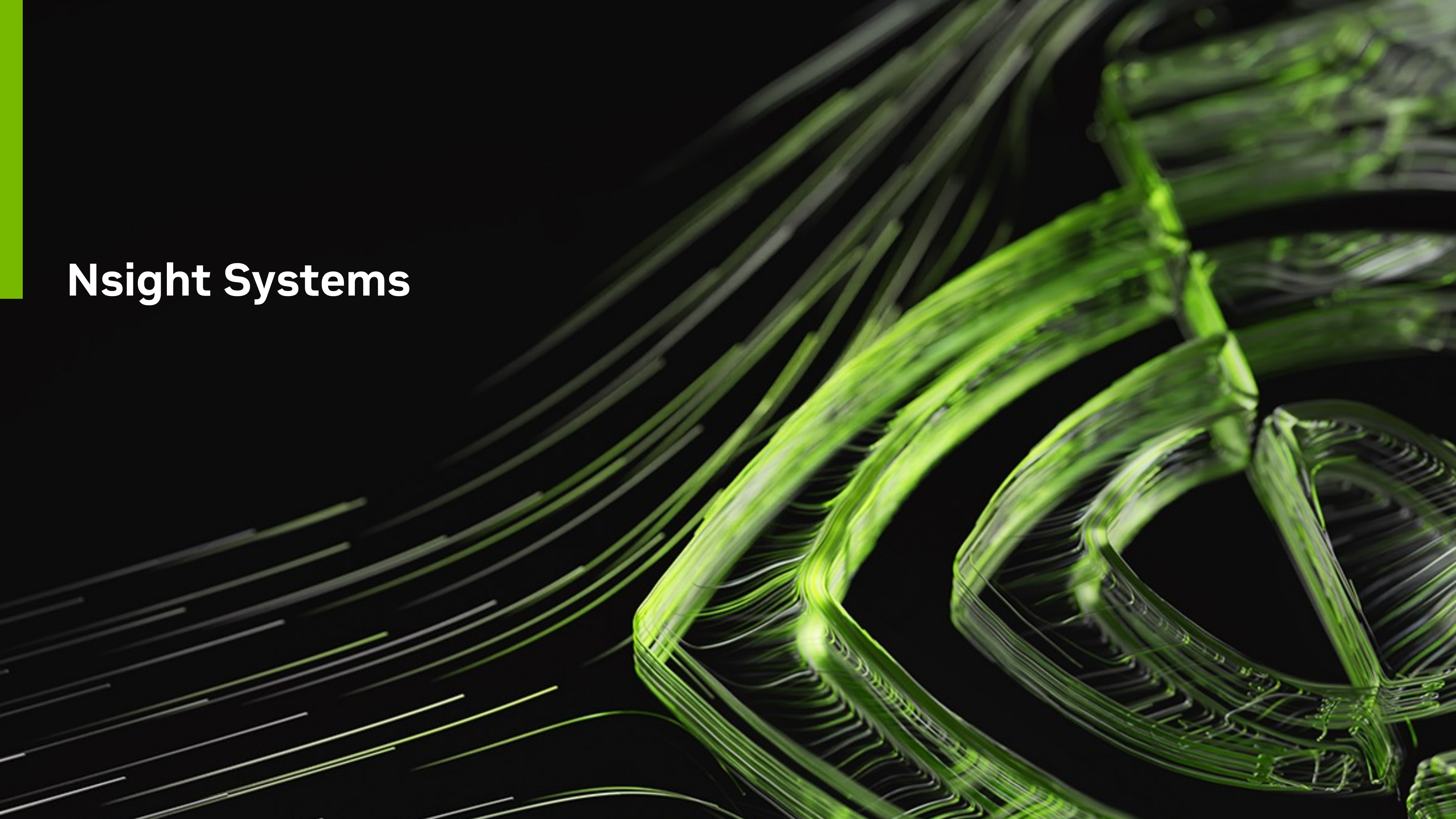
Compute Debuggers/IDEs

New Features

Correctness Tools Features

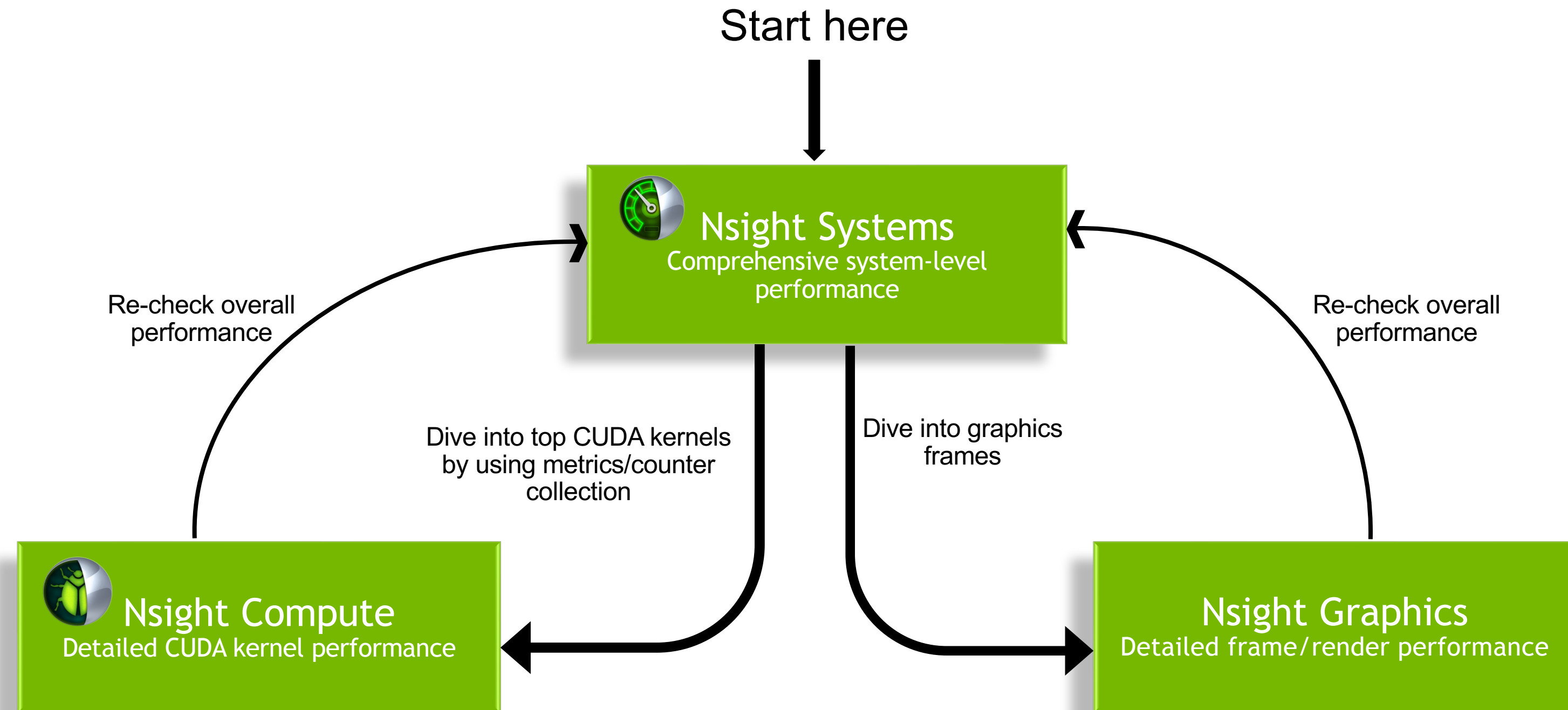
- Debuggers using **Unified Backend**
 - Consistent debugging experience across platforms
 - Performance and usability improvements
 - New features “float all boats”
- New Compute Sanitizer use cases
 - Racecheck support for async-copy, DSMEM (Hopper), and GMMA (Hopper)
 - Memcheck for cache control operations
 - Improved core dump support
 - Multiple OptiX analysis improvements
- Nsight Visual Studio Code Edition
 - New support for debugging containers

```
===== COMPUTE-SANITIZER
===== Invalid __global__ write of size 1 bytes
=====          at 0x4d70 in
/home/cuda/optixBasic/draw_solid_color.cu:69:__raygen__draw_solid_color_0xebf766b2f0642d4e
=====          by thread (0,0,0) in block (0,0,0)
=====          Address 0x7f878f900403 is out of bounds
=====          and is 262,132 bytes after the nearest allocation at
0x7f878f8c0400 of size 16 bytes
=====          Device Frame:NVIDIA internal [0x430]
=====          Saved host backtrace up to driver entry point at
kernel launch time
=====          Host Frame: [0x60fbaa]
=====                          in /lib/x86_64-linux-gnu/libnvoptix.so.1
=====          Host Frame:optix_stubs.h:568:optixLaunch [0xe1ff]
=====                          in /home/cuda/optixBasic/optixBasic
=====          Host
Frame:/home/cuda/optixBasic/optixBasic.cpp:227:main [0xb735]
=====                          in /home/cuda/optixBasic/optixBasic
=====          Host
Frame:../sysdeps/nptl/libc_start_call_main.h:58:__libc_start_call_main
in [0x2dfd0]
=====                          in /lib/x86_64-linux-gnu/libc.so.6
=====          Host Frame:../csu/libc-start.c:379:__libc_start_main
[0x2e07d]
=====                          in /lib/x86_64-linux-gnu/libc.so.6
=====          Host Frame: [0x8dde]
=====                          in /home/cuda/optixBasic/optixBasic
```


The background of the slide is a dynamic, abstract composition. It features a multitude of thin, elongated streaks in various shades of green and white, set against a solid black background. These streaks are oriented diagonally, creating a sense of movement and depth. Some streaks are sharp and bright, while others are blurred, suggesting a long-exposure or motion-capture effect. The overall aesthetic is modern and technological.

Nsight Systems

NSIGHT PROFILERS WORKFLOW



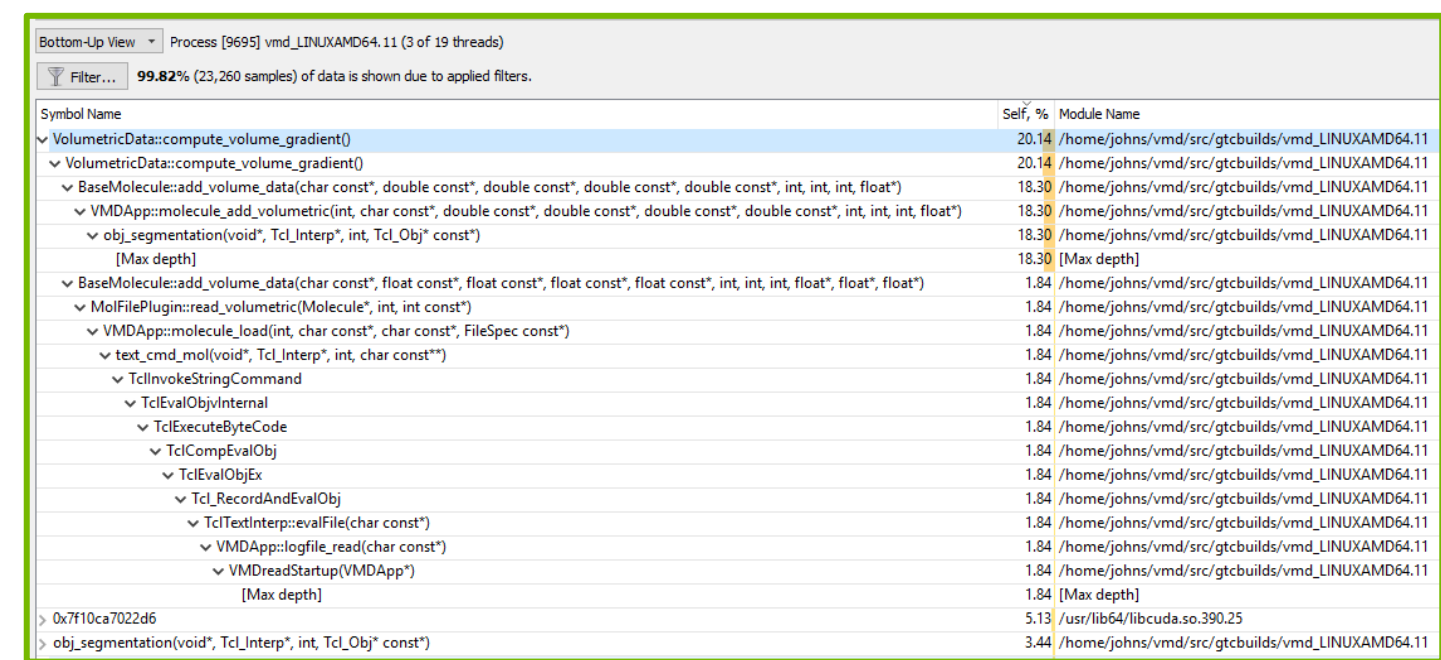
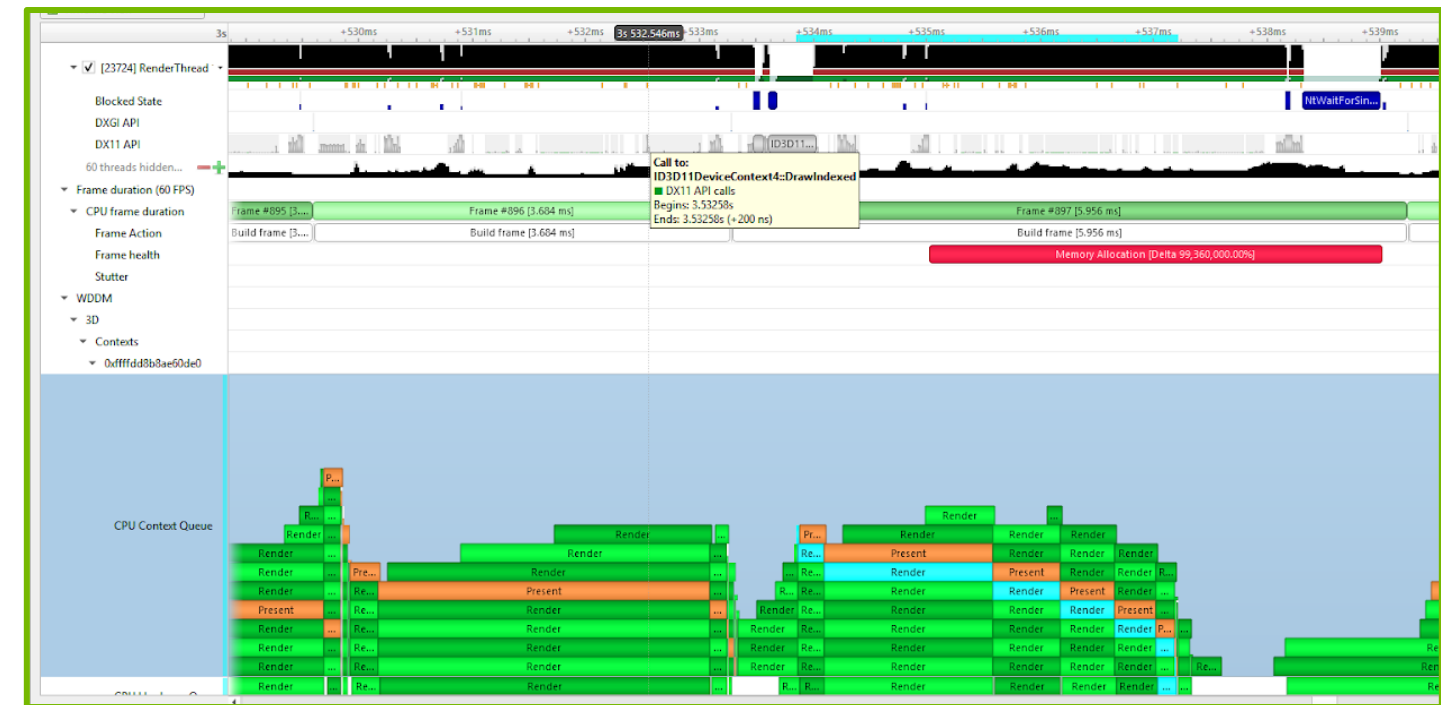
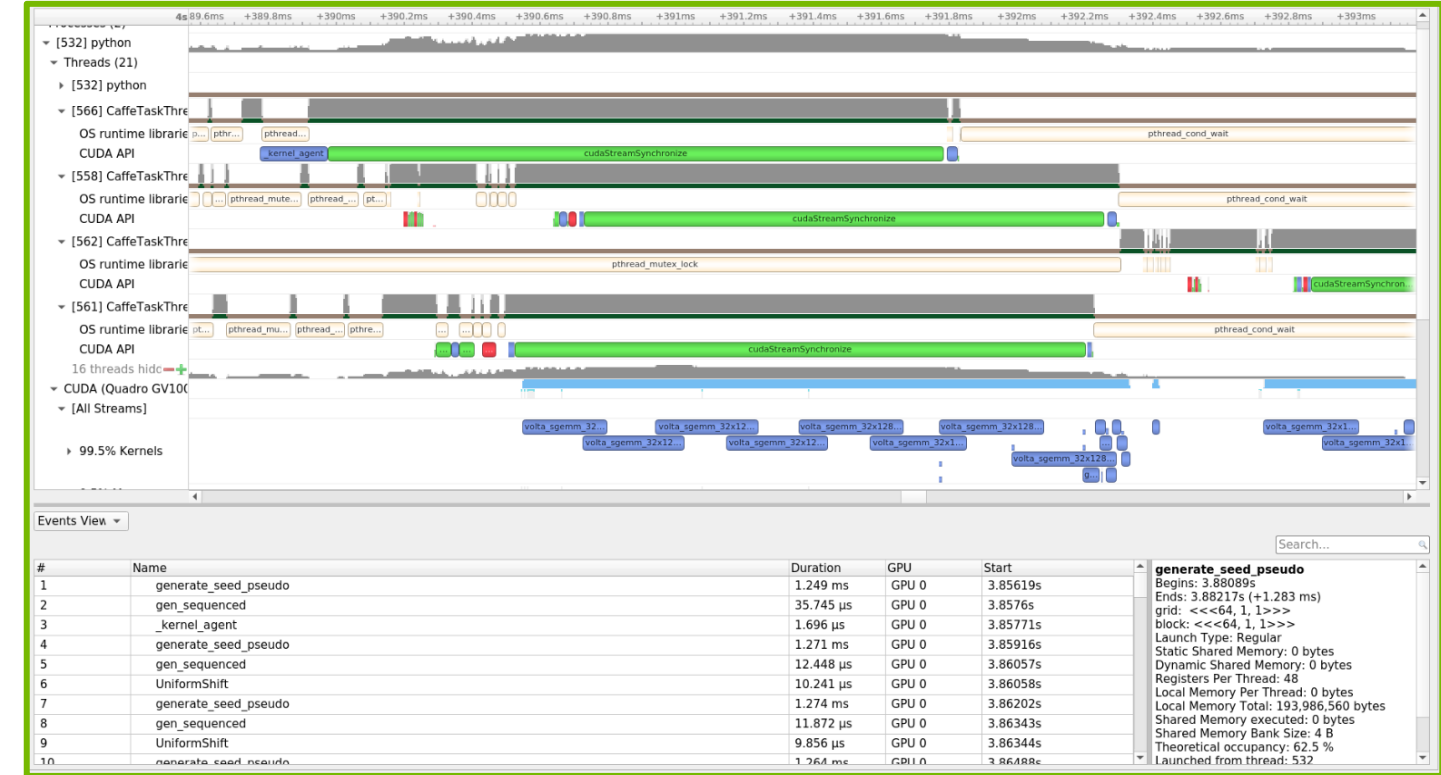


Nsight Systems

System Profiler

Key Features:

- System-wide application algorithm tuning
 - Multi-process tree support
- Locate optimization opportunities
 - Visualize millions of events on a very fast GUI timeline
 - Identify gaps of unused CPU and GPU time
- Balance your workload across multiple CPUs and GPUs
 - CPU algorithms, utilization and thread state
 - GPU streams, kernels, memory transfers, etc
- Command Line, Standalone, IDE Integration
- OS: Linux (x86, Power, ARM Server, Tegra), Windows, macOS X (host)
- GPUs: Pascal+
- Docs/product: <https://developer.nvidia.com/nsight-systems>



CPU cores

GPU metrics
samples

SMs active

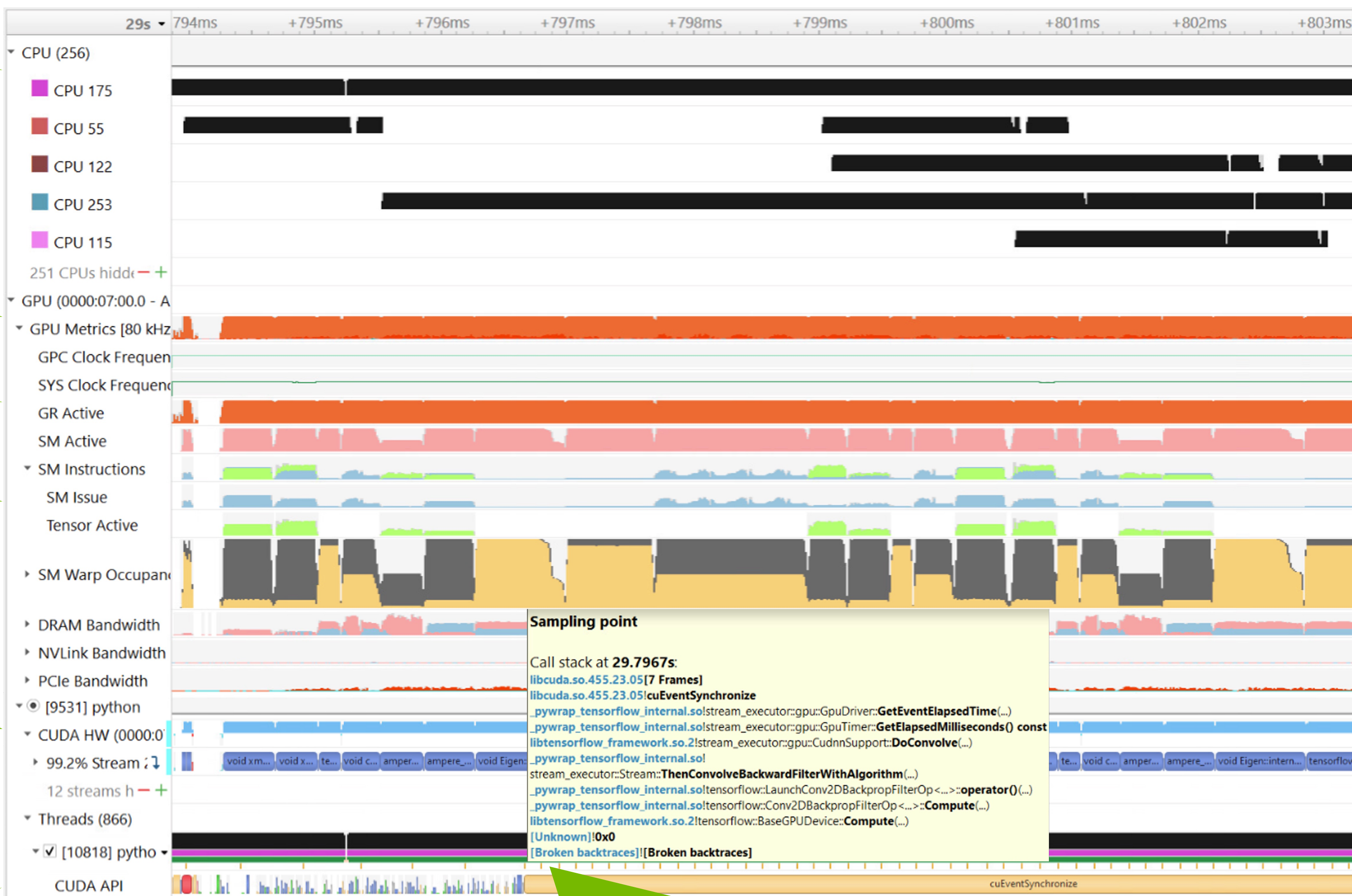
SM instructions
TensorCores

GPU bandwidths

GPU CUDA
kernel & memory
activities

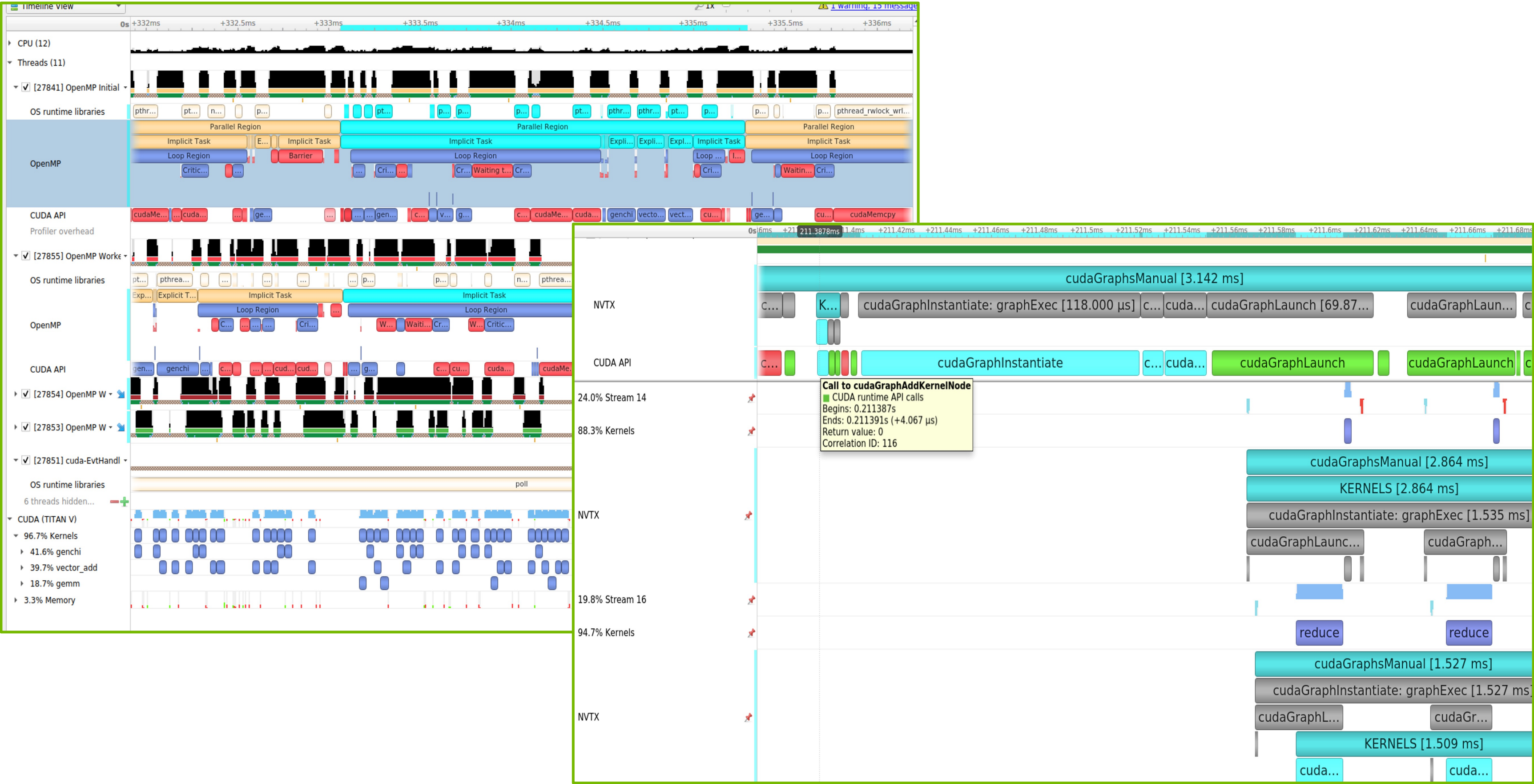
CPU processes
& thread states

CUDA API trace



CPU call-stack samples

Zoom/Filter to Exact Areas of Interest



Pixel Time Coverage Based Level Of Detail (LOD)

Zooming in reveals gaps where there were valleys



Application Profiles with Nsight Systems

```
$ nsys profile -o report -stats=true ./myapp.exe
```

- Generated file: report.qdrep (or report.nsys-rep)
Open for viewing in the Nsight Systems UI
- When using MPI, recommended to use *nsys* after *mpirun/srun*:

```
$ mpirun -n 4 nsys profile ./myapp.exe
```

Core Strategy

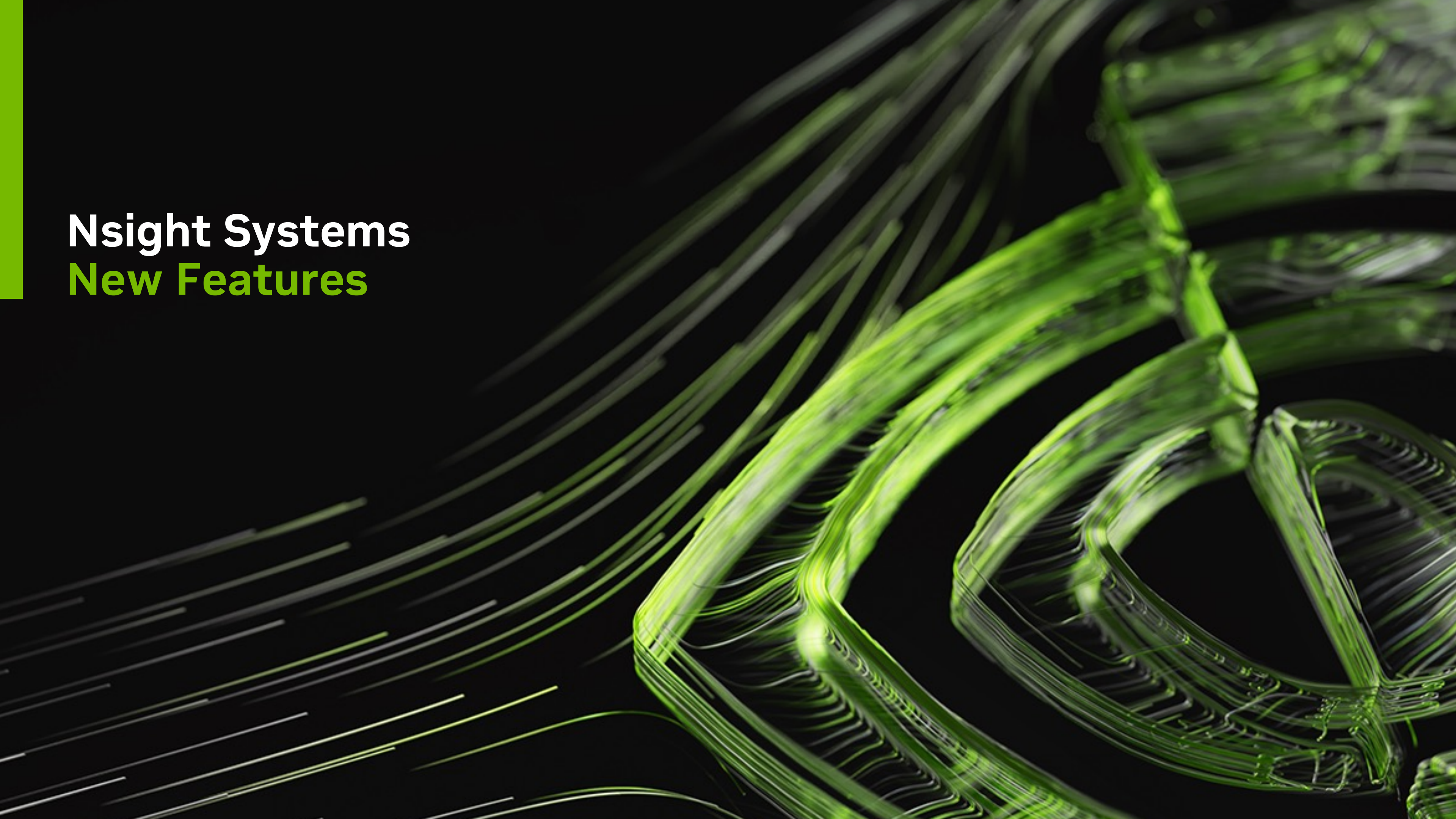
- What's HOT?
 - Will it be easier to shrink what I coded?
 - This is where MOST people concentrate. ...intuitive but not always best !/\$
- What's COLD?
 - Will it be easier to take advantage of the something unused?
 - Free money? Yes please!
- Hot spots might be:
 - Parallelizable?
 - Shrinkable without compromising accuracy, memory, etc
- Cold spots are clear, measurable opportunities!!!
 - How can i remove or fill them?
 - Where do I have incorrect/unnecessary/unexpected dependencies & synchronization?
 - Between threads, processes or across nodes (ex: MPI_Barrier)!

General optimization tips

- Using tensor cores?
- Increase grid and batch size to utilize GPUs width
- Conventional parallelism - more worker threads!
- Parallel pipelining
 - No data dependency? Parallelize!
 - Prefetch next batch/iteration during computation

General optimization tips

- Fuse tiny kernels, copies, memsets.
 - Check out CUDA Graphs
- Overlap/oversubscribe with MPS
- Multi-buffering
 - Don't make everyone wait on the same piece of memory
 - Double, triple buffer
- Avoid moving data back to the CPU
 - Pre-allocate and recycle!
- Minimize managed memory page faults
 - Prefetch!



Nsight Systems **New Features**

CPU Counters and OS Metrics



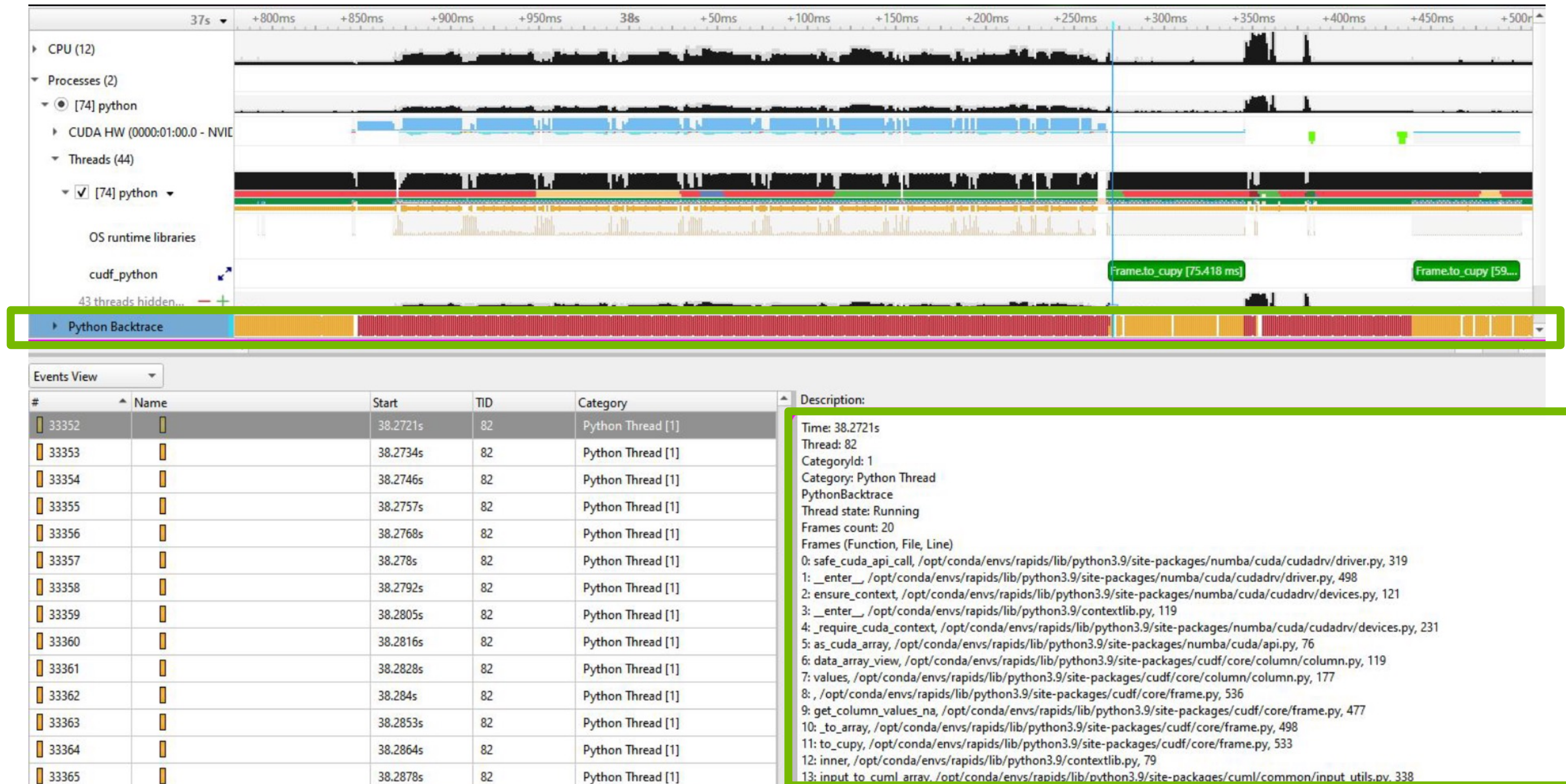
- CPU X86_64/SBSA hardware counters
 - Cache misses, DRAM access, etc...



Python Call Stack Sampling



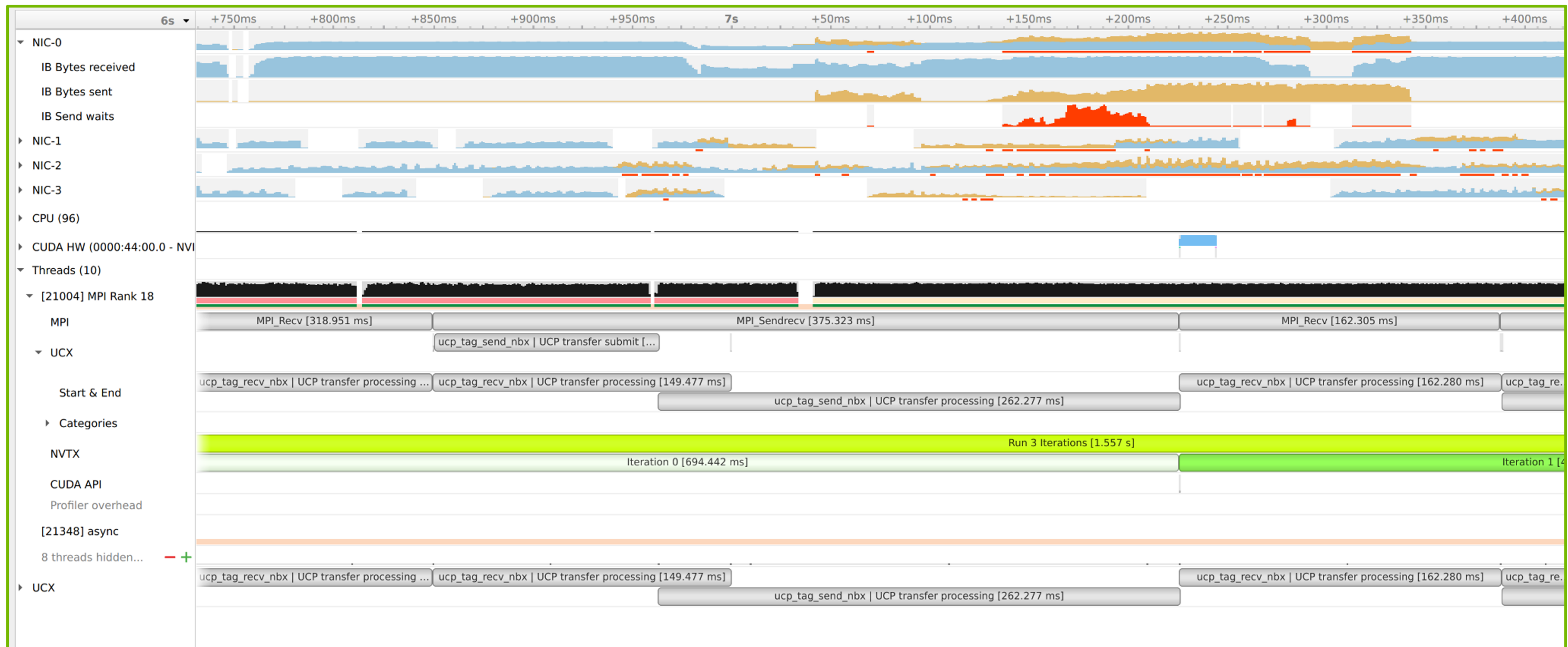
- Periodically sampled stacks
- Correlated with all other collected activity



NVIDIA Networking NIC/DPU Metrics Sampling



- Includes InfiniBand and Ethernet
- Sent / Received / Waits
- Correlate with expected network traffic and other system activities



Network API Trace - Extended Descriptions/Tooltips



- Supports MPI versions
 - OpenMPI
 - MPICH
- Data for each API call
- Identify performance trends and specific underperforming calls

The screenshot displays the 'Events View' of a network API trace tool. It features a list of MPI calls on the left and a detailed description panel on the right. The 'MPI_Bcast' call is selected and highlighted in the list. The description panel provides timing, thread, and data transfer details for this specific call.

Name	Description:
MPI_Init	
MPI_Recv	
MPI_Bcast	MPI_Bcast Begins: 0,164935s Ends: 0,165045s (+109,210 μ s) Thread: 1342434 Bytes sent: 0 Bytes received: 4 Root: 0 MPI_COMM_WORLD
MPI_Bcast	
MPI_Recv	
MPI_Finalize	

Nsight Systems Multi-Node Analysis

Preview in Nsight Systems 2023.2



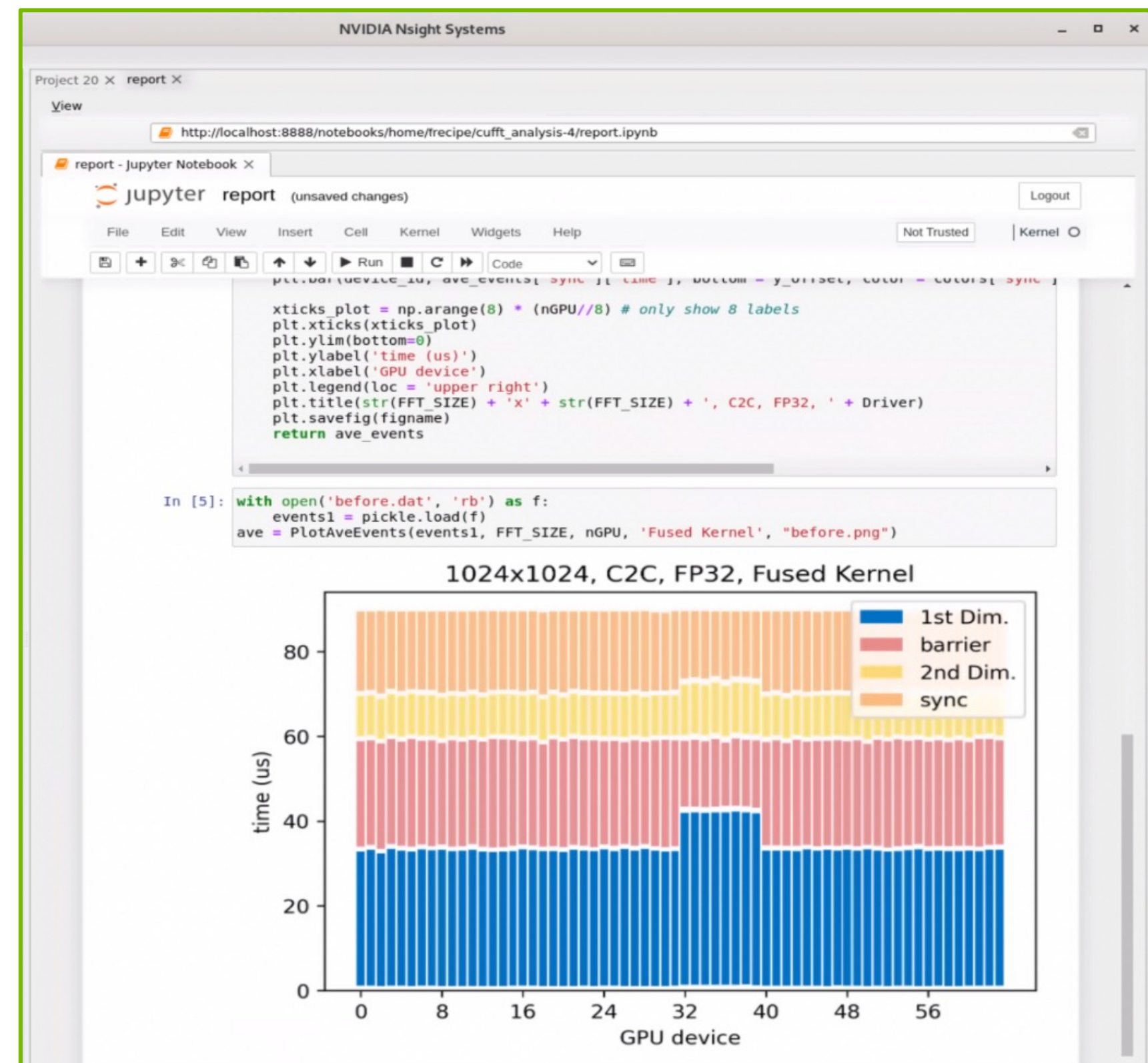
- Recipes to analyze existing reports collected from multiple sources/nodes
- Several default recipes included with Nsight Systems
- Recipes built on extensible Python library
- Recipes generate results various formats:
 - **Jupyter Notebook**
 - Parquet, CSV, ...
- Open result in Nsight Systems GUI or Jupyter Lab

```
> ./nsys recipe -help
```

```
...
cuda_api_sum -- CUDA API Summary
cuda_api_sync -- CUDA Synchronization APIs
cuda_gpu_kern_sum -- CUDA GPU Kernel Summary
cuda_gpu_mem_size_sum -- CUDA GPU MemOps Summary (by Size)
cuda_gpu_mem_time_sum -- CUDA GPU MemOps Summary (by Time)
cuda_memcpy_async -- CUDA Async Memcpy with Pageable Memory
cuda_memcpy_sync -- CUDA Synchronous Memcpy
cuda_memset_sync -- CUDA Synchronous Memset
dx12_mem_ops -- DX12 Memory Operations
gpu_gaps -- GPU Gaps
gpu_time_util -- GPU Time Utilization
nvtx_gpu_proj_trace -- NVTX GPU Trace
nvtx_sum -- NVTX Range Summary
osrt_sum -- OS Runtime Summary
...
```

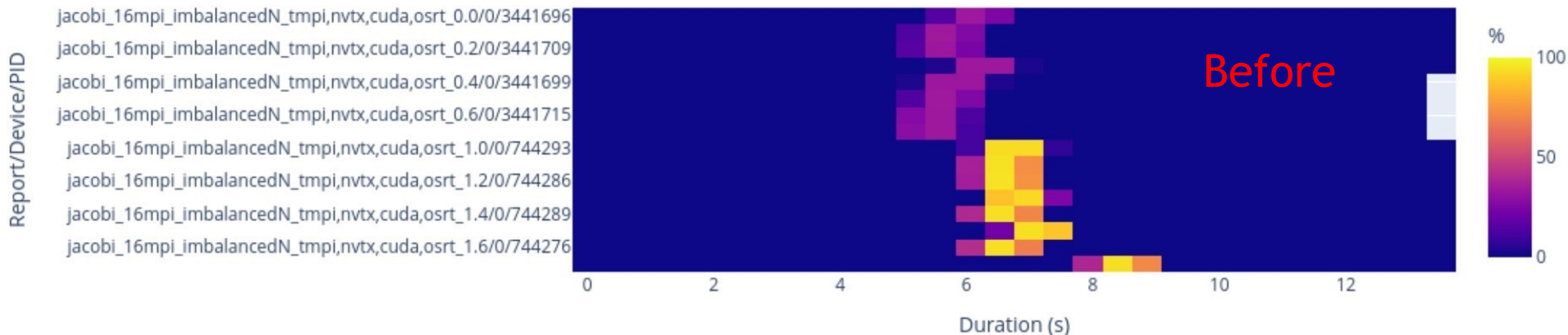
```
> ./nsys recipe cuda_gpu_kern_sum --dir ~/profiles
```

```
> Generated cuda_gpu_kern_sum-1
```

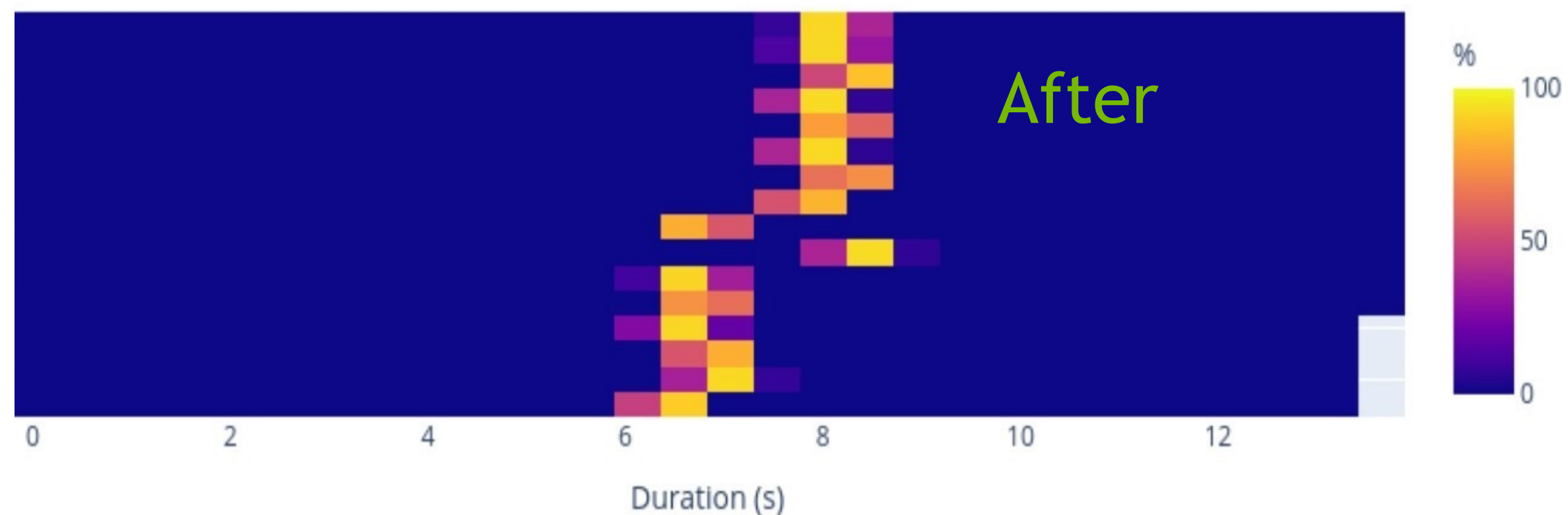


Nsight Systems Multi-Node Analysis

Preview in Nsight Systems 2023.2



- cuda_gpu_time_util_map recipe
- Identify misbehaving GPUs, ranks, or nodes
- Time correlation with application phases
- Click to drill into specific result file for more details



The background features a complex pattern of glowing green lines and shapes on a black field. On the left, numerous thin, parallel green lines radiate outwards. On the right, there are larger, more intricate structures resembling stylized leaves or overlapping rectangular frames, all composed of many fine green lines that create a sense of depth and movement.

NVTX

Tools Extension API

NVIDIA Tools eXtension

NVTX v3

- Decorate application source code with annotations (markers, ranges, nested ranges, ...) to help visualize execution with debugging, tracing and profiling tools

- Header-only library <https://github.com/NVIDIA/NVTX/tree/release-v3/c>.

```
#include <nvtx3/nvToolsExt.h>
```

- Marker:

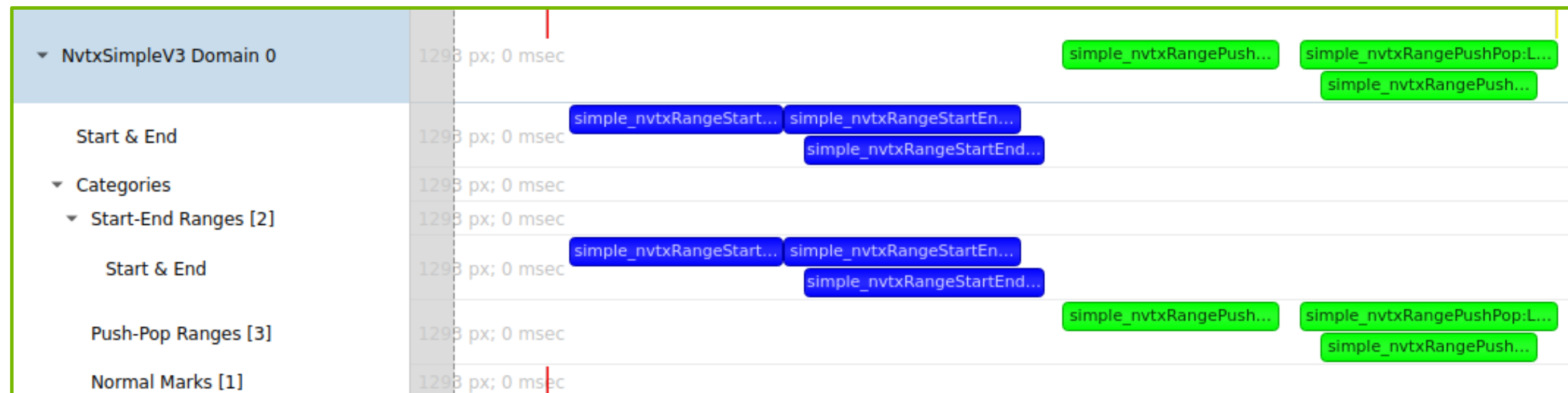
```
nvtxMark("This is a marker");
```

- Push-Pop range

```
nvtxRangePush("This is a push/pop range");  
// Do something interesting in the range.  
nvtxRangePop(); // Pop must be on same thread as corresponding Push
```

- Start-End range

```
nvtxRangeHandle_t handle = nvtxRangeStart("This is a start/end range");  
// Somewhere else in the code, not necessarily same thread as Start call:  
nvtxRangeEnd(handle);
```



API references <https://nvidia.github.io/NVTX/doxygen/index.html> and <https://nvidia.github.io/NVTX/doxygen-cpp/index.html>

NVIDIA SDKs and NVTX

Advanced profiling and performance visualization

GXF

Math Libraries

cuSPARSE

cuSOLVER

cuBLAS

Deep Learning Libraries

DeepStream

TensorRT

cuDNN

cuDLA

Comm. Libraries

NVSHMEM

NCCL

RAPIDS

cuDF

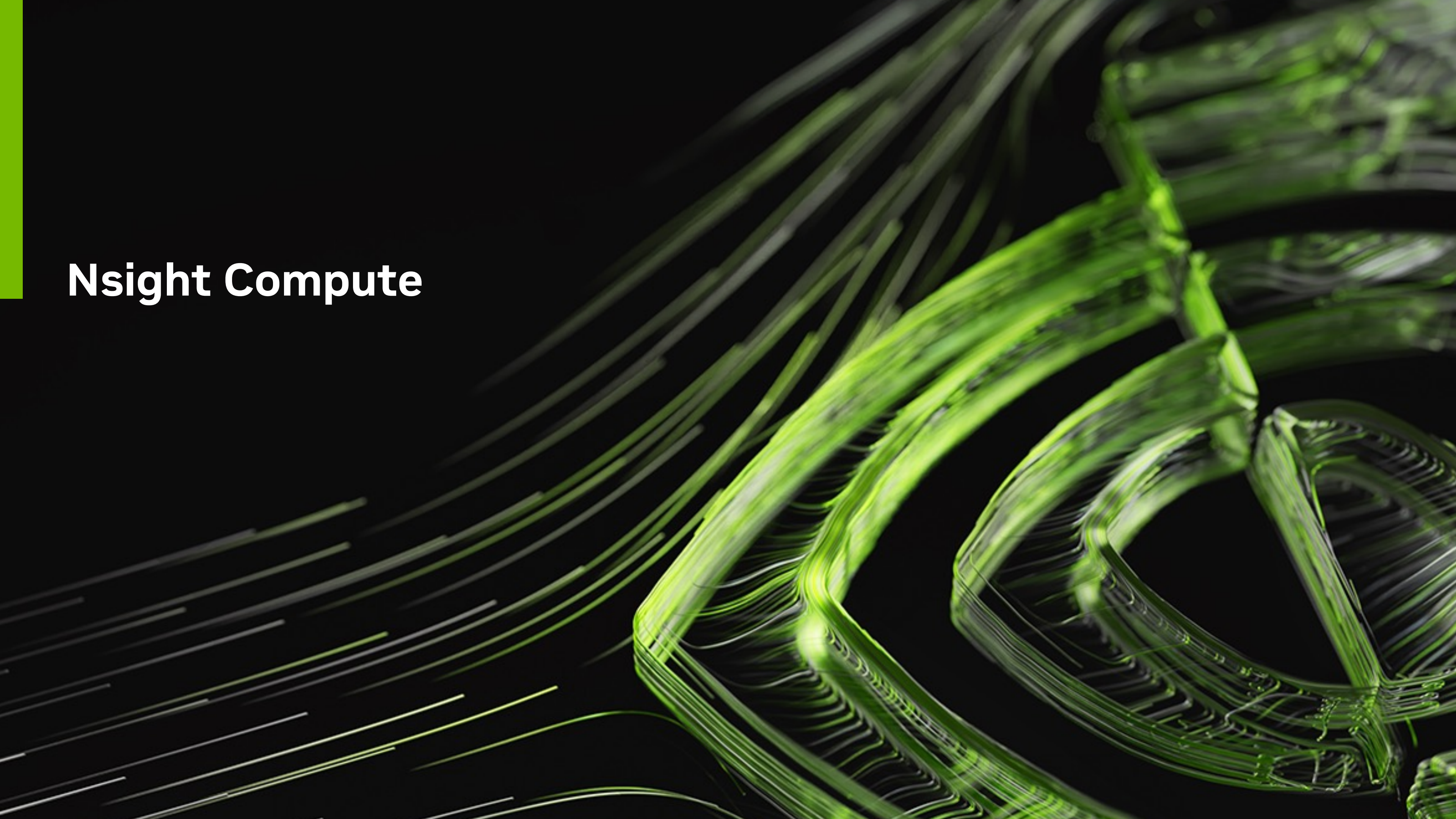
cuFile

cuML

UCX

Profiling DL Models

- Pytorch
 - DNN Layer annotations are disabled by default
 - ++ *"with torch.autograd.profiler.emit_nvtx():"*
 - Manually with *torch.cuda.nvtx.range_(push/pop)*
 - TensorRT backend is already annotated
- Tensorflow
 - Annotated by default with NVTX in NVIDIA TF containers
 - `TF_DISABLE_NVTX_RANGES=1` to disable for production

The background features a complex pattern of glowing green lines and shapes against a solid black field. On the left, numerous thin, parallel green lines radiate outwards. On the right, there are larger, more intricate structures resembling stylized, overlapping leaves or petals, composed of many fine green lines that create a sense of depth and movement.

Nsight Compute

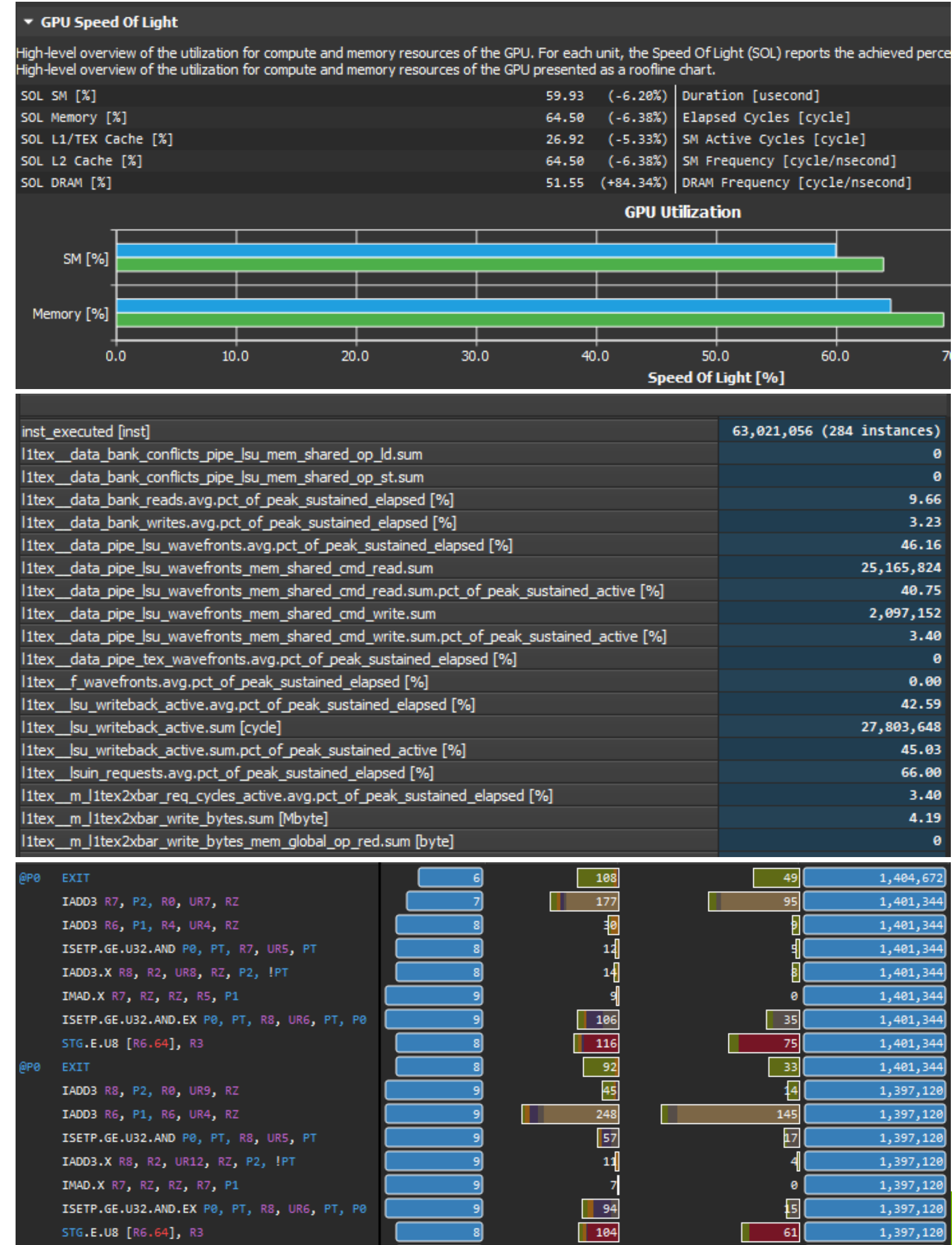


Nsight Compute

Kernel Profiler

Key Features:

- Interactive CUDA API debugging and kernel profiling
- Built-in rules expertise
- Fully customizable data collection and display
- Command Line, Standalone, IDE Integration, Remote Targets
- OS: Linux (x86, Power, Tegra, Arm SBSA), Windows, macOS X (host only)
- GPUs: Volta+
- Docs/product: <https://developer.nvidia.com/nsight-compute>

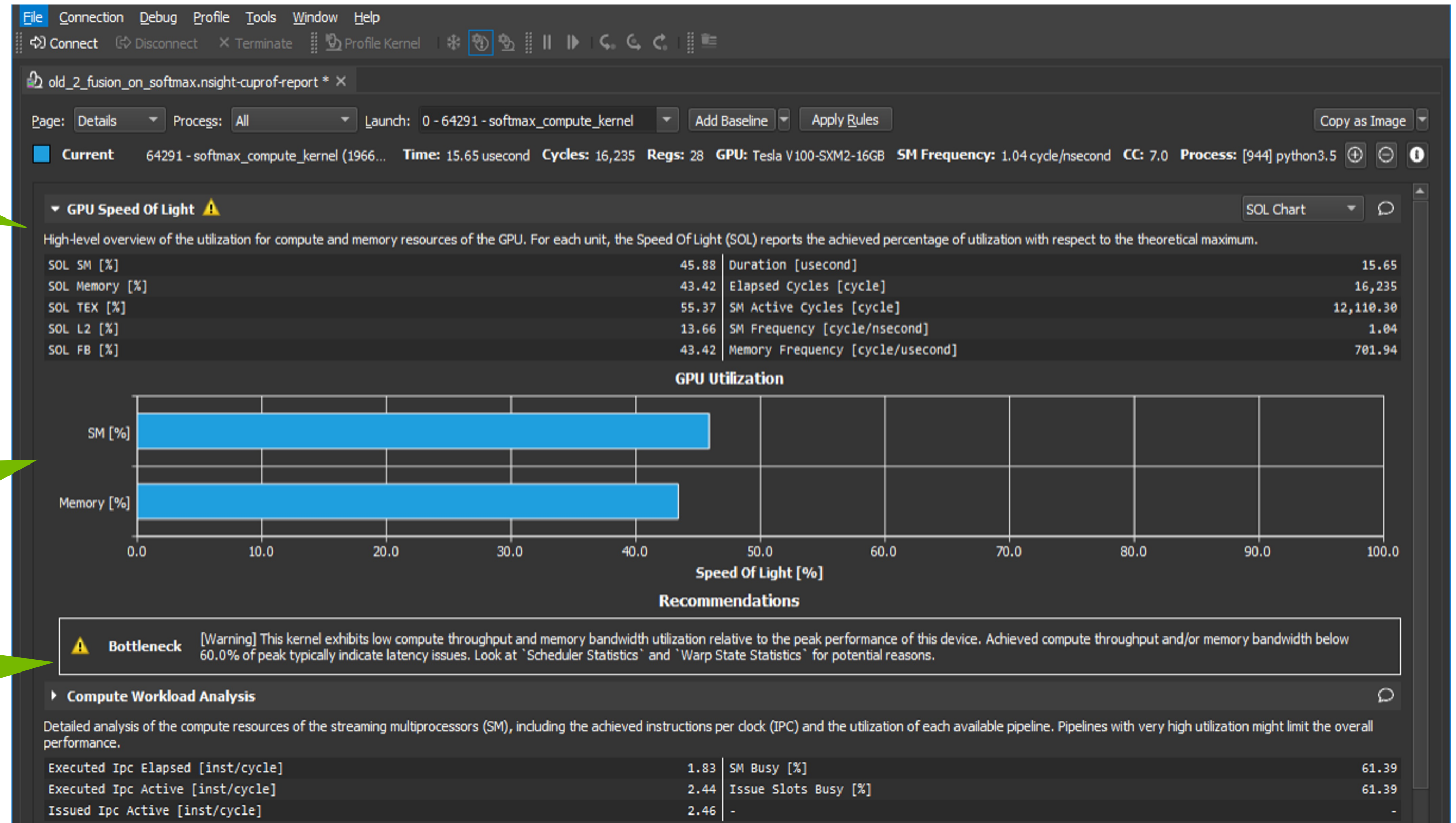


Nsight Compute GUI Interface

Targeted metric sections

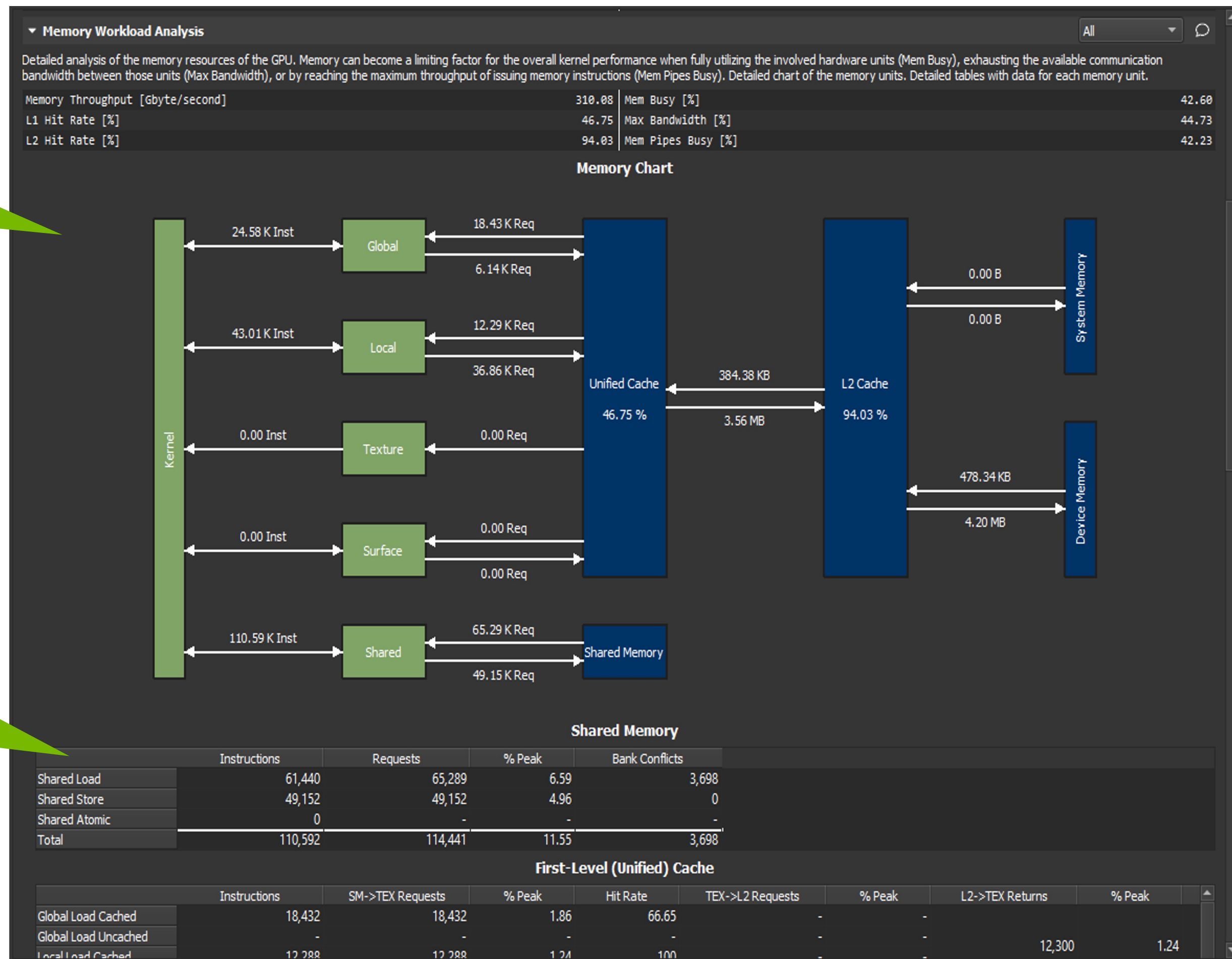
Customizable data collection and presentation

Built-in expertise for Guided Analysis and optimization

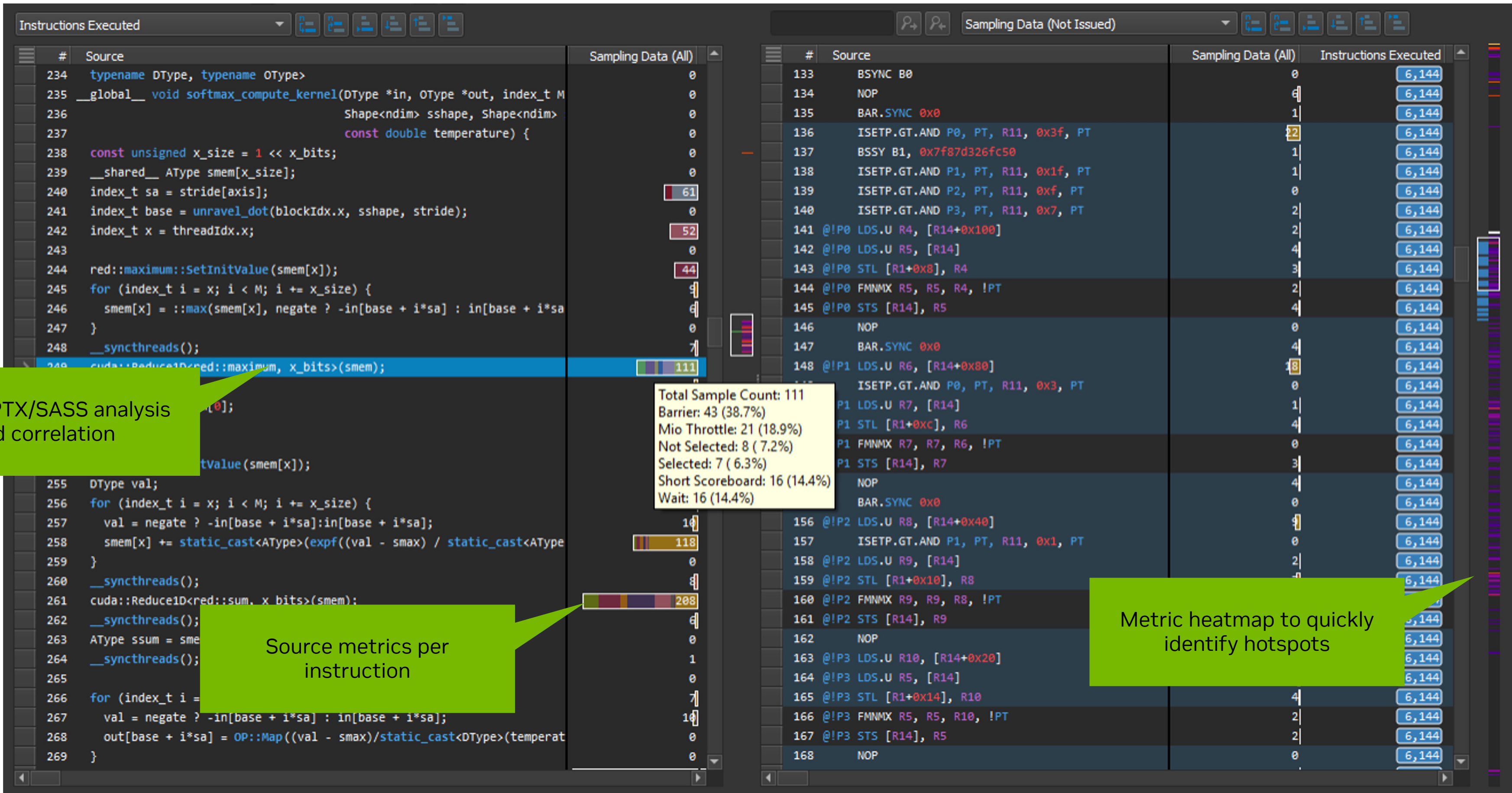


Visual memory analysis chart

Metrics for peak performance ratios



Source/PTX/SASS analysis
and correlation



Kernel profiles with nsight compute

```
$ ncu -k mykernel -o report ./myapp.exe
```

- Generated file: report.ncu-rep
 - Open for viewing in the Nsight Compute UI
- (Without the -k option, Nsight Compute will profile everything and take a long time)

The background is a black field filled with dynamic, glowing green elements. On the left, numerous thin, parallel green lines of varying lengths and orientations create a sense of motion and depth. On the right, there are more complex, overlapping green structures that resemble stylized, translucent leaves or perhaps a network of fibers. These structures have a layered, almost crystalline appearance with some internal detail visible. The overall effect is one of high-tech, organic complexity.

Nsight Compute

New Features

Integrated Basic Trace from Nsight Systems



- Use “System Trace” activity in the connection dialog
- Identify long kernels or compute-bound bottlenecks
- Right-click kernel in timeline to quickly launch profile
- Nsight Compute automatically filters to selected kernel

Project Explorer

Search project...

nsys+ncu

Report 1.nsys-re

Report 1 X

Timeline View

0s +2.6ms +388ms +390ms +392ms +394ms

CUDA HW (0000:2d:00.0 - NVIDIA)

>99.9% Kernels

100.0% MatrixMulCUDA

void ... void MatrixMulCUDA... void MatrixMulCUDA... void MatrixMulCUDA... void MatrixMulCUDA... void MatrixMulCUDA... void MatrixMulCUDA...

Memory

1081 px; 0 msec

1081 px; 0 msec

57] matrixMul

runtime libraries

1081 px; 0 msec

CUDA API

Profiler overhead

1081 px; 0 msec

Profile Kernel

Copy ToolTip

Copy Current Time

Remove Filter

Fit to screen

Undo Zoom (5)

Reset Zoom

Backspace

Ctrl+P

File Connection Debug Profile Tools Window Help

Connect Disconnect Terminate Profile Kernel

Baselines Metric Details

Project Explorer

Search project...

nsys+ncu

Report 1.nsys-rep

report1.ncu-rep

Report 1 X report1.ncu-rep * X

Page: Details Result: 0 - 15 - MatrixMulCUDA

Add Baseline Apply Rules Occupancy Calculator

Result	Time	Cycles	Regs	GPU	SM Frequency
Current	15 - MatrixMulCUDA (20, 10, 1)x(32, 32, 1)	1.53 msecond	20,86,097	31	0 - NVIDIA GeForce RTX 2080 Ti 1.36 cycle/nsecond

GPU Speed Of Light Throughput

High-level overview of the throughput for compute and memory resources of the GPU. For each unit, the throughput reports the achieved percentage of utilization for each individual sub-metric of Compute and Memory to clearly identify the highest contributor.

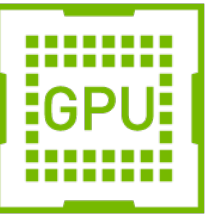
Compute (SM) Throughput [%]	68.45	Duration [msecond]
Memory Throughput [%]	6.38	Elapsed Cycles [cycle]
L1/TEX Cache Throughput [%]	8.38	SM Active Cycles [cycle]
L2 Cache Throughput [%]	0.57	SM Frequency [cycle/nsecond]
DRAM Throughput [%]	0.29	DRAM Frequency [cycle/nsecond]

High Compute Throughput

Compute is more heavily utilized than Memory: Look at the [Compute Workload Analysis](#) section to see what the compute is redundant and could be reduced or moved to look-up tables.

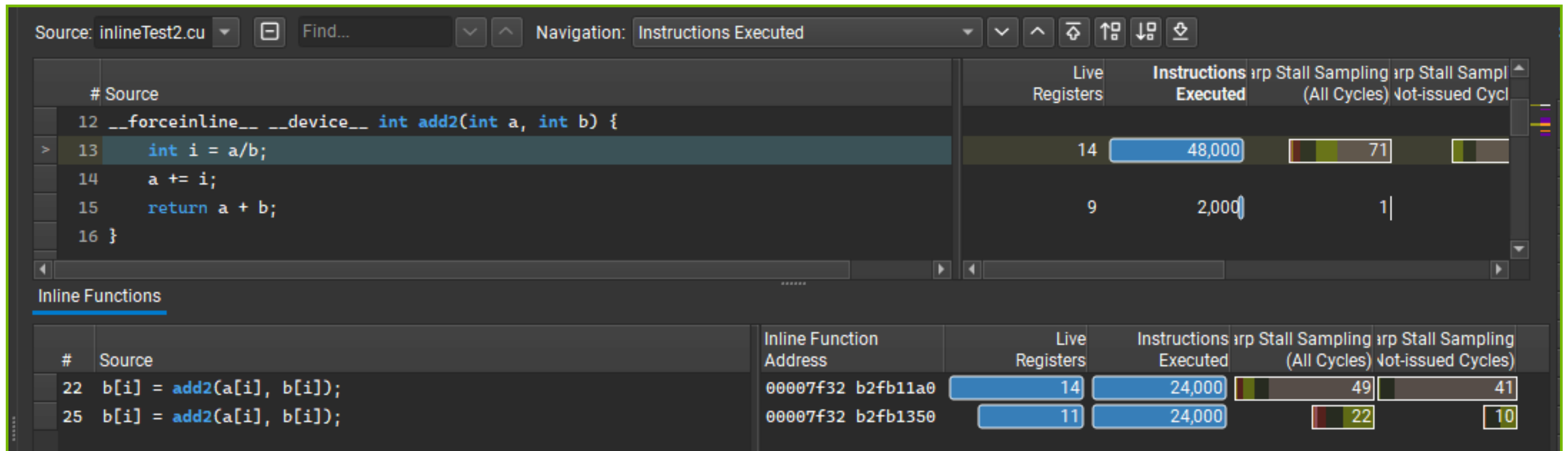
Memory Workload Analysis

Detailed analysis of the memory resources of the GPU. Memory can become a limiting factor for the overall kernel performance when fully utilizing the involved bandwidth between the units (Max Bandwidth) or by pushing the maximum throughput of issuing memory instructions (Max Pipe Busy). Detailed about...



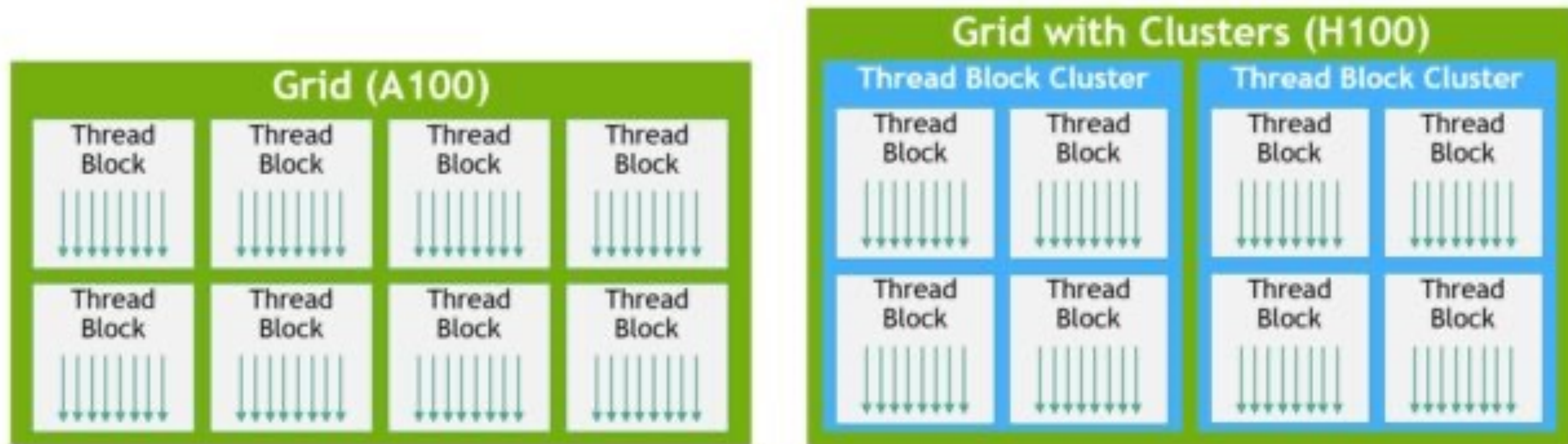
Source Table for Inline Functions

- Metrics can be analyzed per inline site or aggregated for the entire function
- Use compiler `--lineinfo` flag to generate symbols

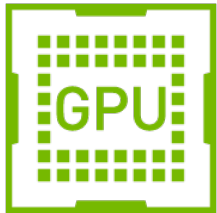


Hopper Thread Block Clusters

- Enables programmatic control of locality at a granularity larger than a single thread block on a single SM
- Clusters enable multiple thread blocks running concurrently across multiple SMs to synchronize and collaboratively fetch and exchange data.
- For more information on Thread Block Clusters see:
 - [CUDA: New Features and Beyond](#)
 - [CUDA Programming Model for Hopper Architecture](#)



Hopper Thread Block Clusters in Nsight Compute



FileConnectionDebugProfileToolsWindowHelp

ConnectDisconnectTerminateProfile KernelBaselines

cask.ncu-rep

Page: DetailsResult: 2 - 26908 - warp_specialized_kernel_2Add BaselineApply RulesOccupancy CalculatorCopy as Image

Current

Result

26908 - warp_specialized_kernel_gmma_utmaldbg (1, 2, 1)x(70, 2, 1)x(384, 1, ...

Time

88,992.00 nsecond

Cycles

142,518

Regs

168

GPU

0 - NVIDIA Graphics Device

SM Frequency

1,597,445,556.90 cycle/second

CC

9.0

Launch Statistics

Cluster Size	2.00	Cluster Scheduling Policy	cudaClusterSchedulingPolicySpread
Grid Size	140.00	Registers Per Thread [register/thread]	168.00
Block Size	384.00	Static Shared Memory Per Block [byte/block]	0.00
Threads [thread]	107,520.00	Dynamic Shared Memory Per Block [byte/block]	218,368.00
Waves Per SM	1.00	Driver Shared Memory Per Block [byte/block]	1,024.00
Function Cache Configuration	cudaFuncCachePreferNone	Shared Memory Configuration Size [byte]	233,472.00

Occupancy

Theoretical Occupancy [%]	18.75	Block Limit Registers [block]	1.00
Theoretical Active Warps per SM [warp]	12.00	Block Limit Shared Mem [block]	1.00
Achieved Occupancy [%]	15.56	Block Limit Warps [block]	5.00
Achieved Active Warps Per SM [warp]	9.96	Block Limit SM [block]	32.00
Max Cluster Size	16.00	Max Active Clusters	70.00

Source Counters

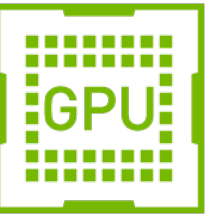
Branch Instructions [inst]	1,576,420.00	Branch Efficiency [%]	99.86
Branch Instructions Ratio [%]	0.05	Avg. Divergent Branches	4.00

Hopper SM 90

Cluster Scheduling Policy

Cluster Launch Size

Cluster Occupancy



New Documented Samples

Working to Ease the Complexity of Performance Optimization

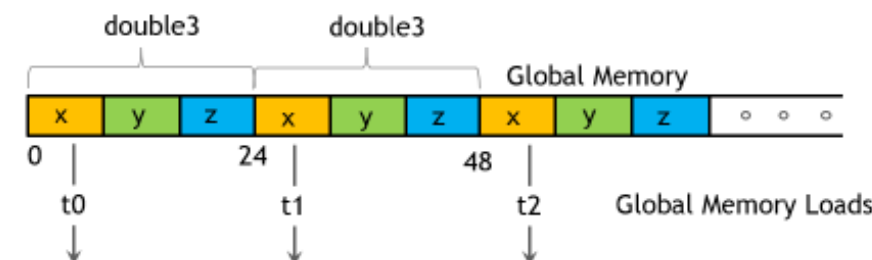
- Packaged with Nsight Compute installation
- Includes reproducible instructions, descriptions, pre-collected results, and source code
- Explaining common performance issues in CUDA applications
- **Uncoalesced Global Accesses** and **Shared Memory Bank Conflicts**

Chapter 4. INITIAL VERSION OF THE KERNEL

The initial version of the sample code provides a naive implementation for the kernel which adds a floating point constant to an input array of double3.

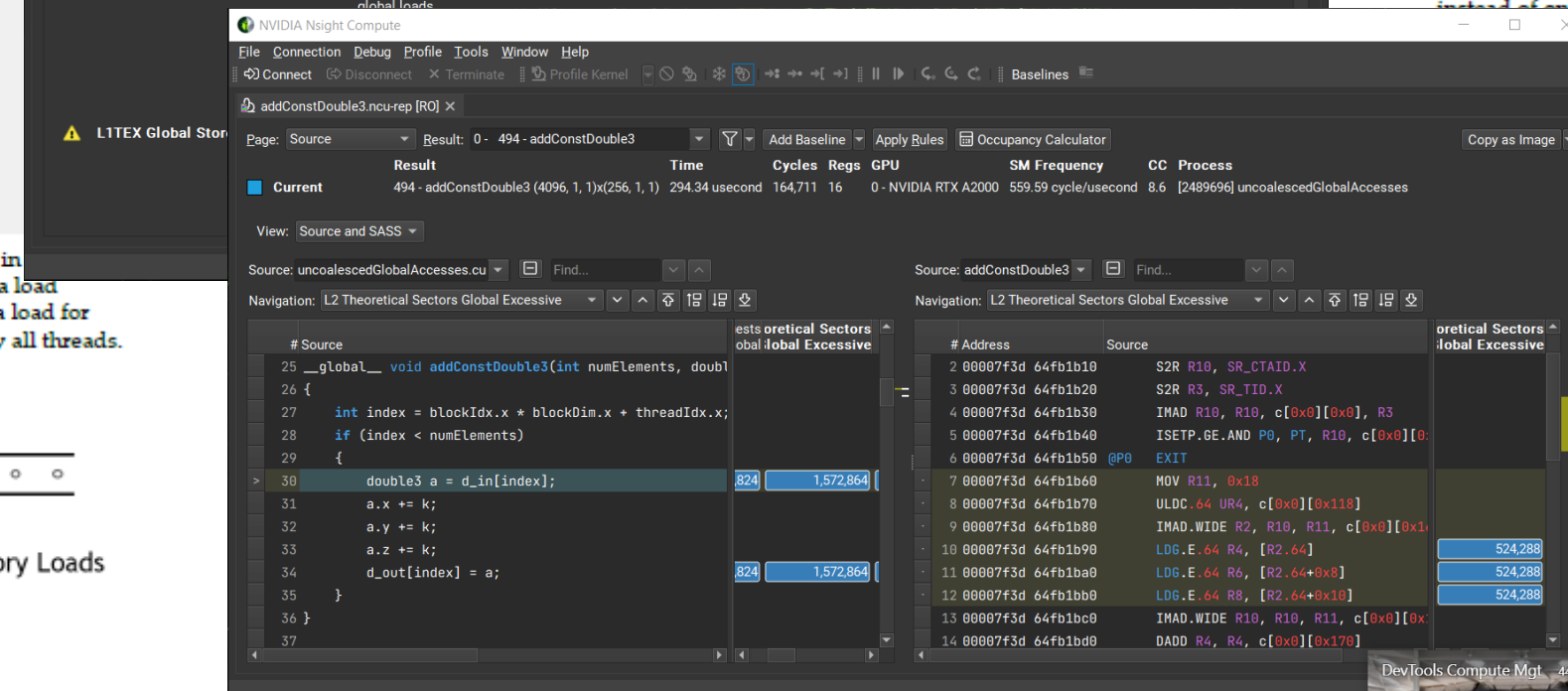
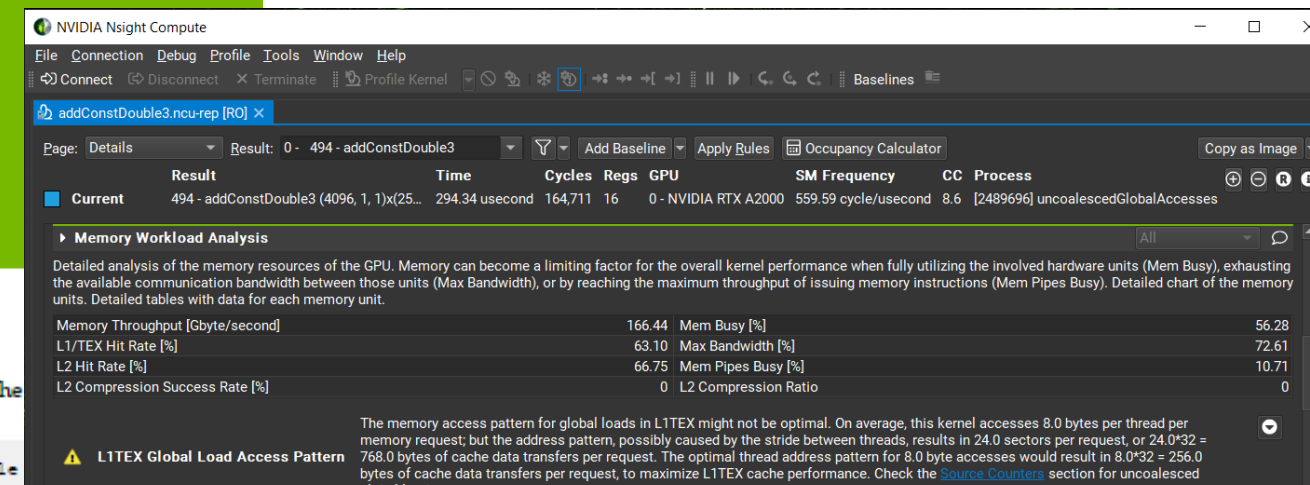
```
__global__ void addConstDouble3(int numElements, double3 *d_in, double3 *d_out)
{
    int index = blockIdx.x * blockDim.x + threadIdx.x;
    if (index < numElements)
    {
        double3 a = d_in[index];
        a.x += k;
        a.y += k;
        a.z += k;
        d_out[index] = a;
    }
}
```

The instruction `a = d_in[index]` in the kernel code results in each thread in accessing global memory 24-bytes apart. In the first step all threads request a load for `d_in[index].x` as shown in the following diagram. In the second step a load for `d_in[index].y` and in the third step a load for `d_in[index].z` is made by all threads.



Profile the initial version of the kernel

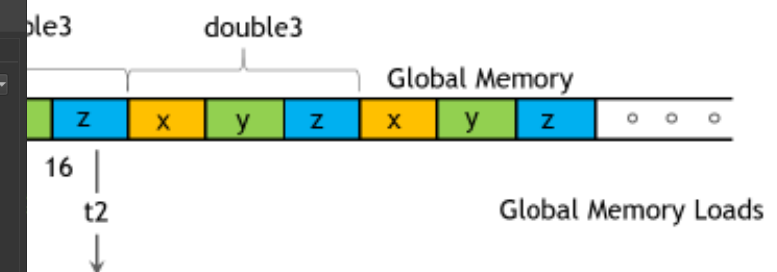
- Refer to the README distributed with the sample on how to build the application
- Run `ncu-ui` on the host system



Chapter 5. UPDATED VERSION OF THE KERNEL

Considering the uncoalesced accesses reported by the profiler we analyze the global load access pattern. Each thread executes 3 reads for the three double values in double3.

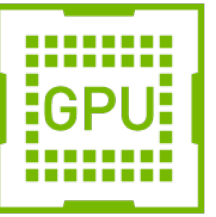
We can treat the double3 array as a double array and each thread can process one double instead of one double3. With this change threads in a warp access consecutive double loads and stores are coalesced.



```
__global__ void addConstDouble(int numElements, double *d_in, double *d_out)
{
    int index = blockIdx.x * blockDim.x + threadIdx.x;
    if (index < numElements)
    {
        double a = d_in[index];
        a += k;
        d_out[index] = a;
    }
}
```

Updated kernel

The execution time has reduced from 294 microseconds to 239 microseconds. We can set the initial version of the kernel and compare the profiling results.



Source Page Guidance

- **Branch Navigation**
 - Quickly jump to target of branch instructions
- **Rules Markers**
 - Automated analysis and rules output per source line
 - Quick detection and navigation to lines that matter
 - Explanation of detected issues

```
10 #include <stdio.h>
11 #include <cuda_runtime_api.h>
12
13 #define BLOCK_SIZE 256
14
15 #define RUNTIME_API_CALL(apiFuncCall)
16 do {
17     cudaError_t _status = apiFuncCall;
18     if (_status != cudaSuccess) {
19         fprintf(stderr, "%s:%d: error: function
20             __FILE__, __LINE__, #apiFuncCal
21         exit(EXIT_FAILURE);
22     }
23 } while (0)
24
25 __global__ void addConstDouble3(int numElements
26 {
27     int index = blockIdx.x * blockDim.x + thre
28     if (index < numElements)
29     {
30         double3 a = d_in[index];
31
32         double3 b = d_in[index + BLOCK_SIZE];
33         double3 c = d_in[index + 2 * BLOCK_SIZE];
34         double3 d = d_in[index + 3 * BLOCK_SIZE];
35         double3 e = d_in[index + 4 * BLOCK_SIZE];
36         double3 f = d_in[index + 5 * BLOCK_SIZE];
37         double3 g = d_in[index + 6 * BLOCK_SIZE];
38         double3 h = d_in[index + 7 * BLOCK_SIZE];
39         double3 i = d_in[index + 8 * BLOCK_SIZE];
40         int index = blockIdx.x * blockDim.x + thre
```

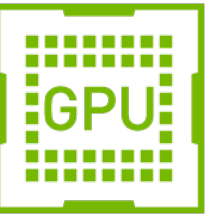
1	0.10%		0.1
3	10.86%		11.6
2	0.48%		0.4
11	30.00%		29.8
9	19.95%		19.6
9	1.23%		1.2
7	0.68%		0.6
9	15.45%		15.5
1	21.25%		21.3

```
11 00007f3d 64fb1ba0
12 00007f3d 64fb1bb0
13 00007f3d 64fb1bc0
14 00007f3d 64fb1bd0
15 00007f3d 64fb1be0
16 00007f3d 64fb1bf0
17 00007f3d 64fb1c00
18 00007f3d 64fb1c10
19 00007f3d 64fb1c20
20 00007f3d 64fb1c30
21 00007f3d 64fb1c40
22 00007f3d 64fb1c50
23 00007f3d 64fb1c60
24 00007f3d 64fb1c70
25 00007f3d 64fb1c80
26 00007f3d 64fb1c90
27 00007f3d 64fb1ca0
28 00007f3d 64fb1cb0
29 00007f3d 64fb1cc0
30 00007f3d 64fb1cd0
31 00007f3d 64fb1ce0
32 00007f3d 64fb1cf0
```

⚠ This line has 524288 excessive sectors from uncoalesced global accesses at SASS address 0x7f3d64fb1b90.

⚠ This line has 524288 excessive sectors from uncoalesced global accesses at SASS address 0x7f3d64fb1ba0.

⚠ This line has 524288 excessive sectors from uncoalesced global accesses at SASS address 0x7f3d64fb1bb0.



Potential Speedup Calculations

- Estimated speedup potential calculated for detected issues
- Sort kernels based on potential speedup
- Helps users understand where to focus effort and potential benefit

ID	Potential Speedup ^	Issues	Function Name	Demangled Name	Duration	Compute Throughput	Memory Throughput	# Registers	Device Name	Process	Grid Size	Block Size
219	18.42	7	bitonicMerge...	bitonicMerge...	44.29	11.19	88.93	22	NVIDIA GeFor...	[1117161] sorting...	2048, 1...	256, 1...
106	16.48	7	bitonicMerge...	bitonicMerge...	43.68	11.36	89.84	22	NVIDIA GeFor...	[1117161] sorting...	2048, 1...	256, 1...
109	16.24	7	bitonicMerge...	bitonicMerge...	42.82	11.27	89.30	22	NVIDIA GeFor...	[1117161] sorting...	2048, 1...	256, 1...
288	15.79	7	bitonicMerge...	bitonicMerge...	42.66	11.48	90.66	22	NVIDIA GeFor...	[1117161] sorting...	2048, 1...	256, 1...
52	15.47	7	bitonicMerge...	bitonicMerge...	42.78	11.28	89.41	22	NVIDIA GeFor...	[1117161] sorting...	2048, 1...	256, 1...
77	15.45	7	bitonicMerge...	bitonicMerge...	43.30	11.43	90.43	22	NVIDIA GeFor...	[1117161] sorting...	2048, 1...	256, 1...
258	15.41	7	bitonicMerge...	bitonicMerge...	44.38	11.23	88.65	22	NVIDIA GeFor...	[1117161] sorting...	2048, 1...	256, 1...
79	15.28	7	bitonicMerge...	bitonicMerge...	43.55	11.21	88.65	22	NVIDIA GeFor...	[1117161] sorting...	2048, 1...	256, 1...
207	14.14	7	bitonicMerge...	bitonicMerge...	42.88	11.24	88.36	22	NVIDIA GeFor...	[1117161] sorting...	2048, 1...	256, 1...
275	14.04	7	bitonicMerge...	bitonicMerge...	42.85	11.12	88.76	22	NVIDIA GeFor...	[1117161] sorting...	2048, 1...	256, 1...
150	14	7	bitonicMerge...	bitonicMerge...	44.96	11.08	88.75	22	NVIDIA GeFor...	[1117161] sorting...	2048, 1...	256, 1...

Discovered performance optimization opportunities for this result. See the Details page for more info.

Issue Slot Utilization Potential Speedup: 16.48x

Every scheduler is capable of issuing one instruction per cycle, but for this kernel each scheduler only issues an instruction every 16.5 cycles. This might leave hardware resources underutilized and may lead to less optimal performance. Out of the maximum of 12 warps per scheduler, this kernel allocates an average of 8.28 active warps per scheduler, but only an average of 0.09 warps were eligible per cycle. Eligible warps are the subset of active warps that are ready to issue their next instruction. Every cycle with no eligible warp results in no instruction being issued and the issue slot remains unused. To increase the number of eligible warps, reduce the time the active warps are stalled by inspecting the top stall reasons on the [Warp State Statistics](#) and [Source Counters](#) sections.

long_scoreboard Potential Speedup: 5.19x

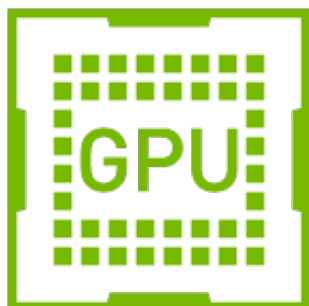
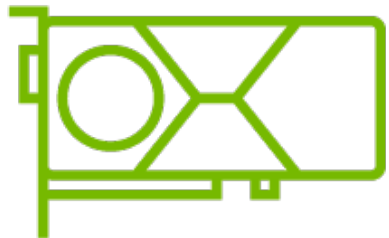
On average, each warp of this kernel spends 110.1 cycles being stalled waiting for a scoreboard dependency on a L1TEX (local, global, surface, texture) operation. Find the instruction producing the data being waited upon to identify the culprit. To reduce the number of cycles waiting on L1TEX data accesses verify the memory access patterns are optimal for the target architecture, attempt to increase cache hit rates by increasing data locality (coalescing), or by changing the cache configuration. Consider moving frequently used data to shared memory.. This stall type represents about 80.7% of the total average of 136.4 cycles between issuing two instructions.

The background is a black field filled with numerous thin, bright green lines that radiate from the left side towards the right. On the right side, there are several larger, more complex green shapes that appear to be composed of many overlapping, slightly curved lines, giving them a three-dimensional, crystalline or fibrous appearance. These shapes are interconnected and seem to be the result of the lines converging and folding back on themselves.

Putting it all together

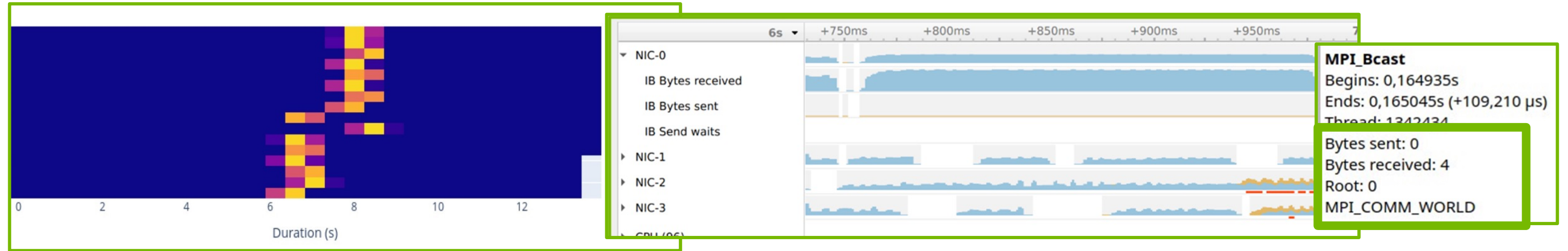
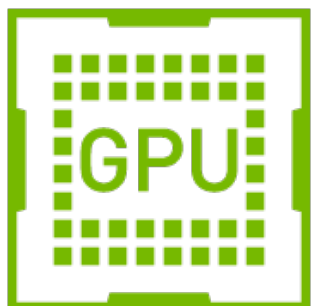
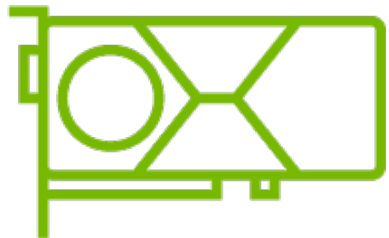
From the Macro to the Micro

Putting it all together



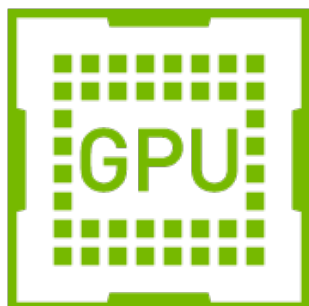
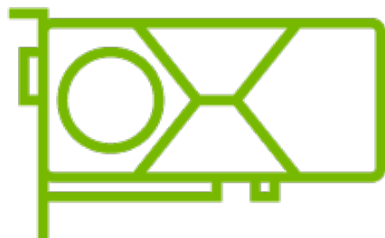
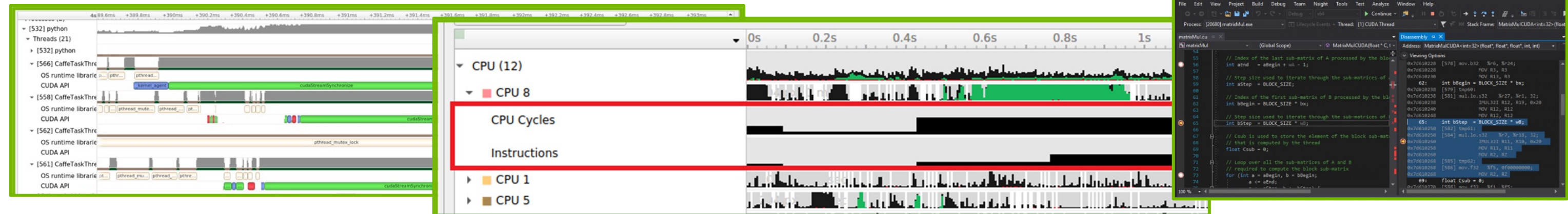
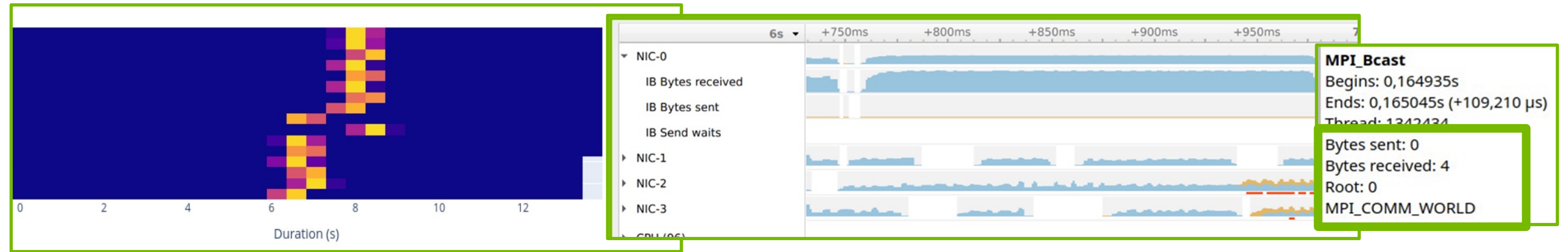
From the Macro to the Micro

Putting it all together



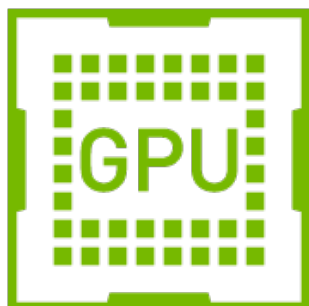
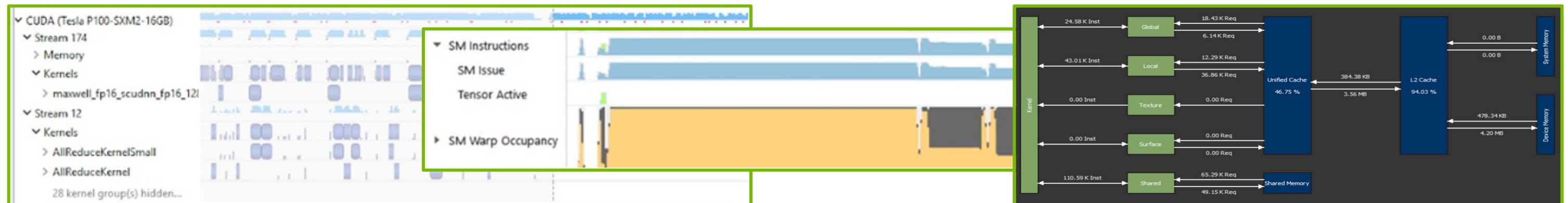
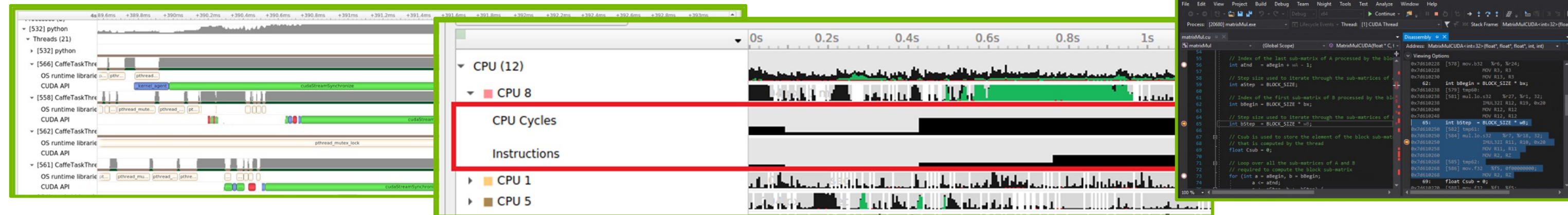
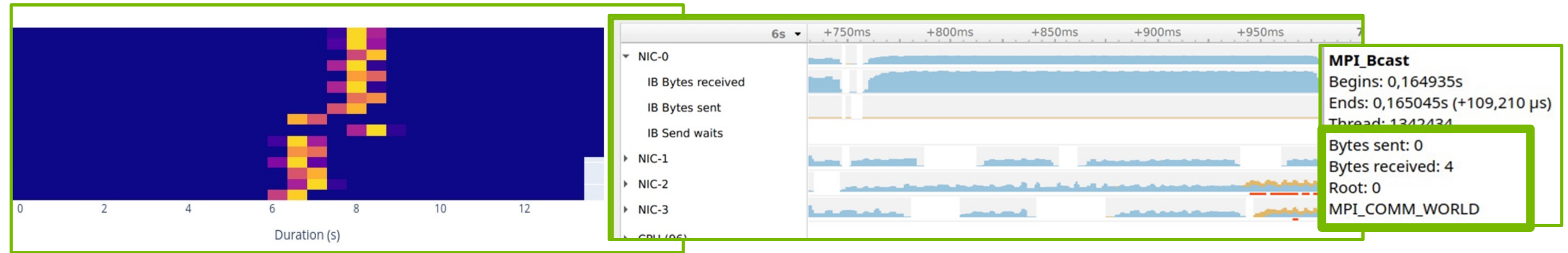
From the Macro to the Micro

Putting it all together



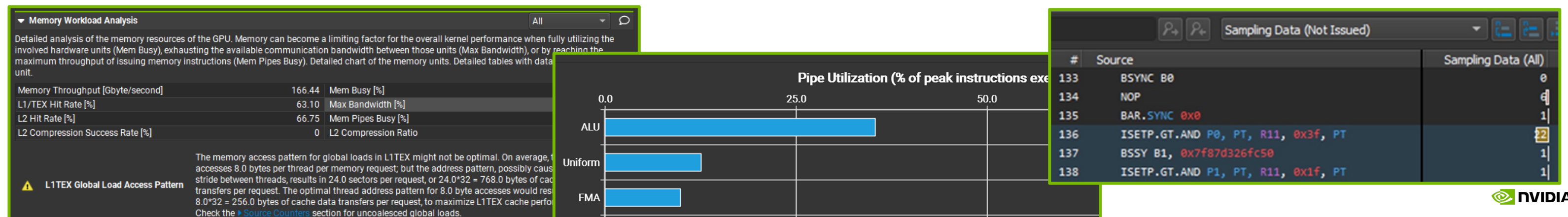
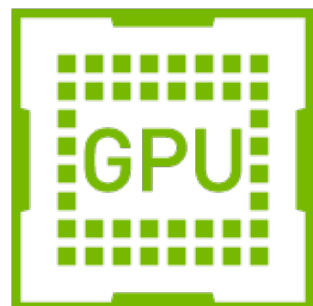
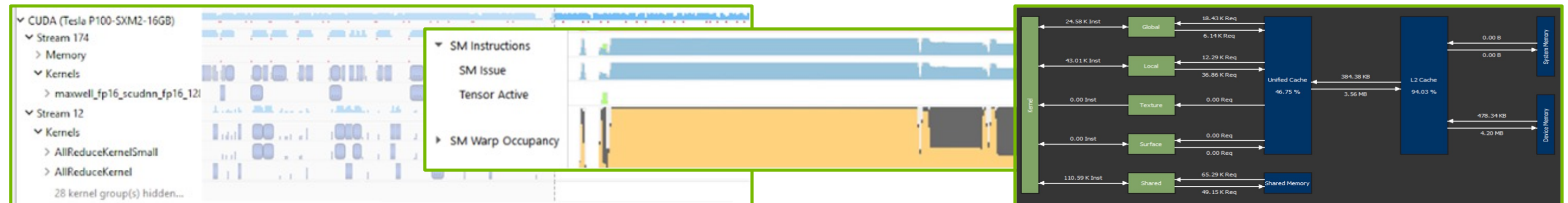
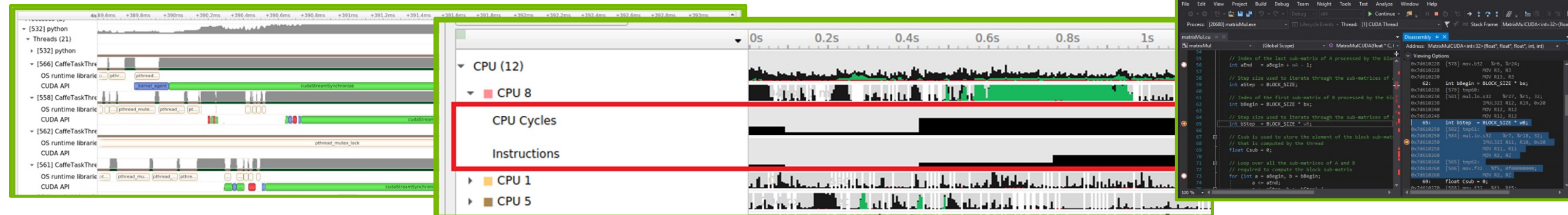
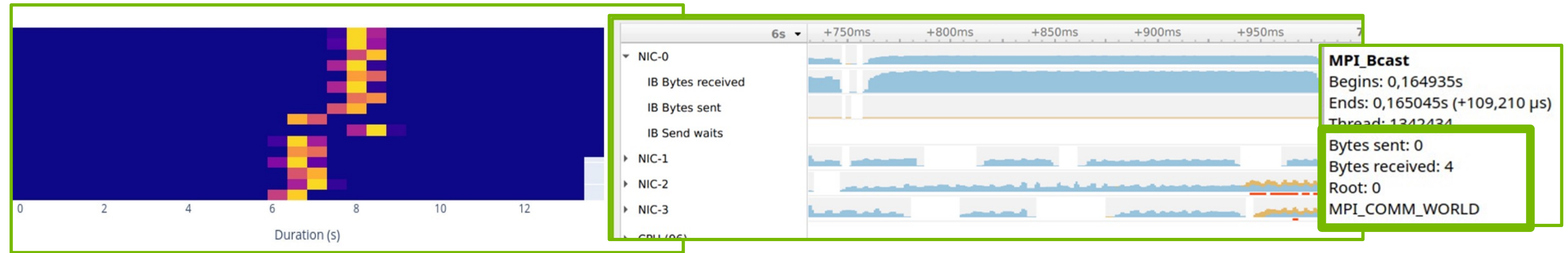
From the Macro to the Micro

Putting it all together



From the Macro to the Micro

Putting it all together



The background of the slide is a black field filled with numerous thin, curved, and slightly blurred lines in shades of green and yellow. These lines create a sense of motion and depth, resembling a microscopic view of fibers or a stylized representation of light trails. On the far left, there is a solid, bright green vertical bar.

Additional Resources

OpenHackathons.org

[Events](#)[Attendees](#)[Mentors](#)[About](#)[OpenACC](#)[Log In](#)

Cyfronet AI for Science Bootcamp

Dates: September 11-12, 2023

Event Focus: AI

Event Format: Virtual Event

Region: Europe/Middle East/Africa

Application Status: [Open](#)

CESGA AI for Science Bootcamp

Dates: September 18-19, 2023

Event Focus: AI

Event Format: Virtual Event

Region: Europe/Middle East/Africa

Application Status: [Open](#)

NASA Open Hackathon

Dates: September 12-21, 2023

Event Focus: HPC+AI

Event Format: Virtual Event

Region: North America/Latin America

Application Status: Closed

CSCS Open Hackathon

Dates: September 18-22, 2023

Event Focus: HPC+AI

Event Format: In-Person Event

Region: Europe/Middle East/Africa

Application Status: Closed

NERSC AI for Scientific Computing Bootcamp

Dates: October 18-20, 2023

Event Focus: AI

Event Format: Virtual Event

Region: North America/Latin America

Application Status: [Open](#)

NSM Open Hackathon

Dates: October 10-20, 2023

Event Focus: HPC+AI

Event Format: Hybrid Event

Region: Asia-Pacific

Application Status: [Open](#)

Developer Tools Across GTC

- Sessions
 - [S51205](#): From the Macro to the Micro - CUDA Developer Tools Find and Fix Problems at Any Scale
 - [S51421](#): Optimizing at Scale: Investigating and Resolving Hidden Bottlenecks for Multi-Node Workloads
 - [S51882](#): Become Faster in Writing Performant CUDA Kernels using the Source Page in Nsight Compute
 - [S51772](#): Debugging CUDA: An Overview of CUDA Correctness Tools
 - [S51230](#): Orin Performance Bodybuilding with Nsight Developer Tools
 - [SE52434](#): Jetson Edge AI Developer Days: Getting the Most Out of Your Jetson Orin Using NVIDIA Nsight Developer Tools
- Labs
 - [DLIT51143](#): Master Common Optimization Patterns Efficiently with Nsight Profiling Tools
 - [DLIT51202](#): Debugging and Analyzing Correctness of CUDA Applications
 - [DLIT51580](#): Ray-Tracing Development using NVIDIA Nsight Graphics and NVIDIA Nsight Systems
- Connect with Experts
 - [CWES52036](#): What's in Your CUDA Toolbox? CUDA Profiling, Optimization, and Debugging Tools for the Latest Architectures
 - [CWES52009](#): Using NVIDIA Developer Tools to Optimize Ray Tracing
- Developer Tools are free and packaged in the latest version of the CUDA Toolkit
 - <https://developer.nvidia.com/cuda-downloads>
- Support is available via:
 - <https://forums.developer.nvidia.com/c/development-tools/>
- More information at:
 - <https://developer.nvidia.com/tools-overview>

