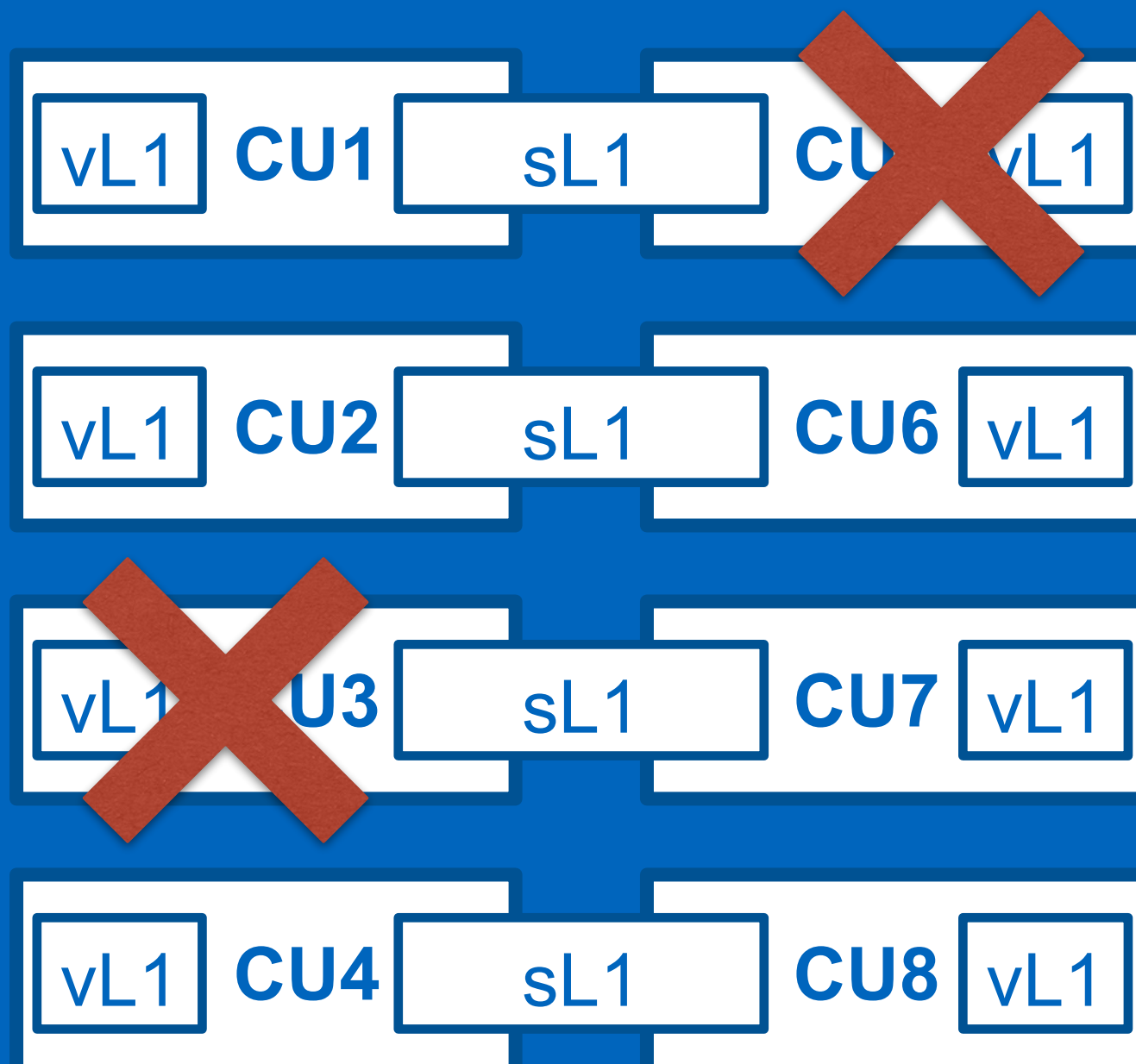


Microbenchmark deep-dive: AMD scalar L1 cache sharing

AMD CDNA architecture



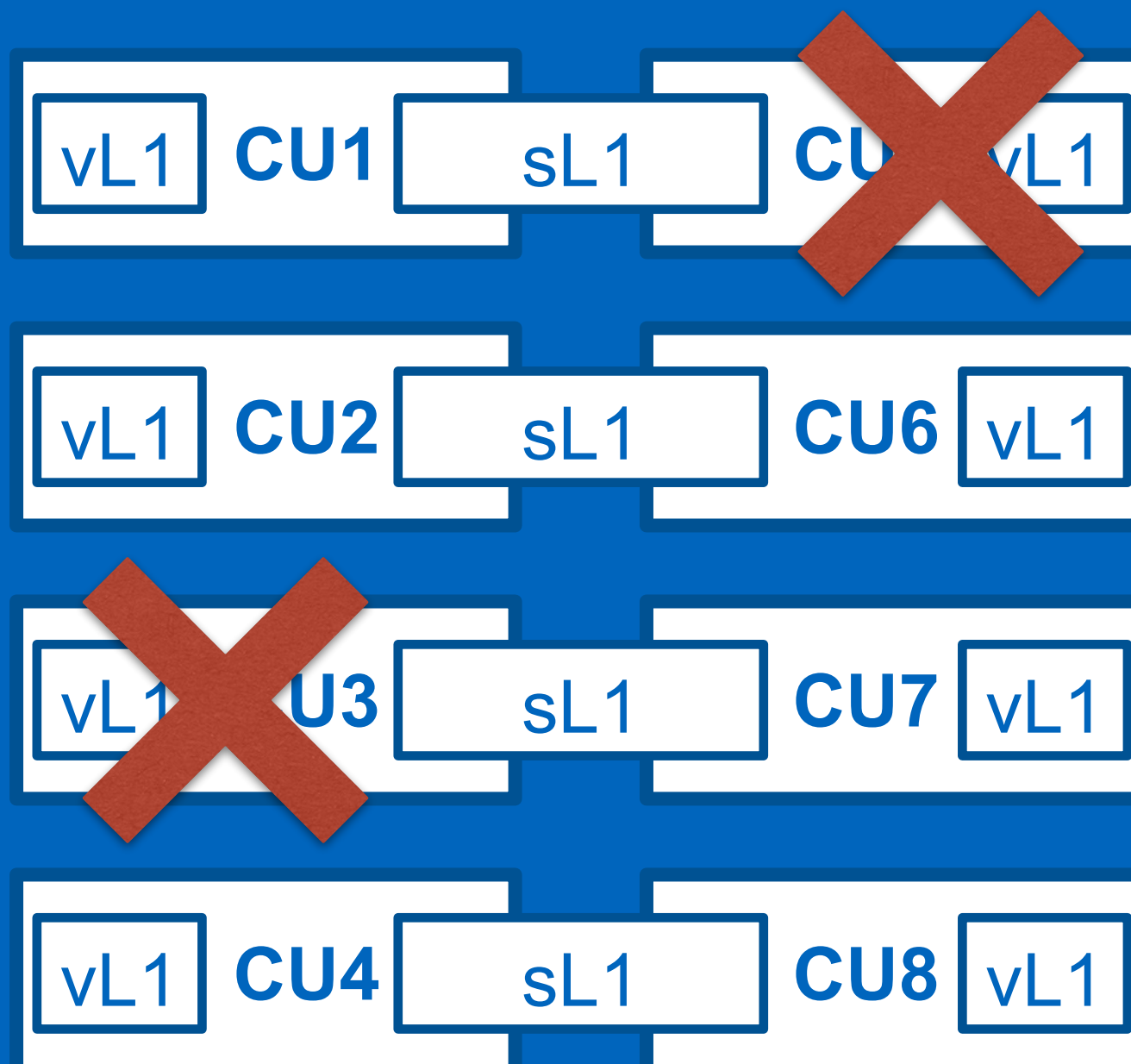
Why is it interesting?

1

How many bytes per CU available?

Microbenchmark deep-dive: AMD scalar L1 cache sharing

AMD CDNA architecture

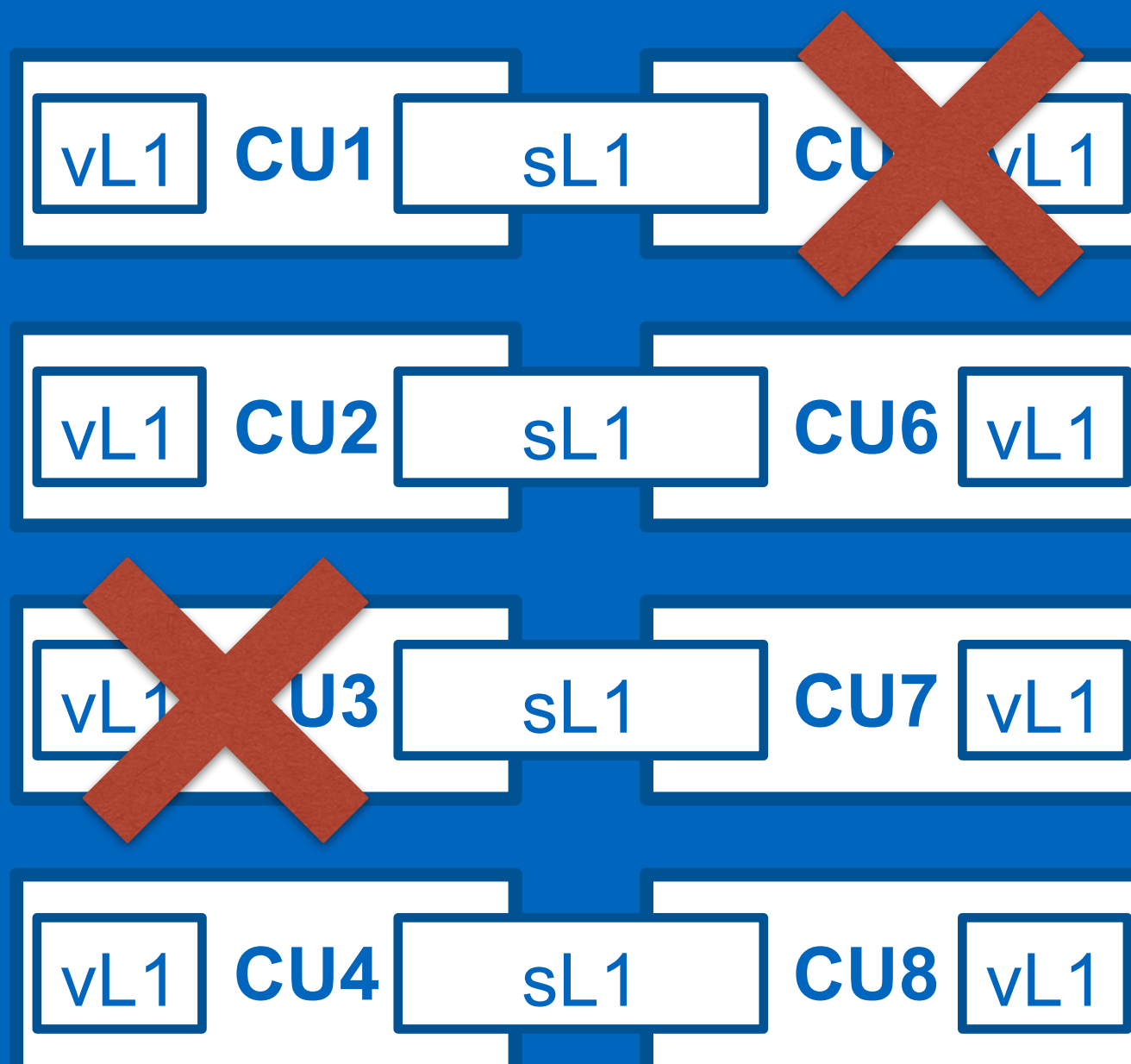


Why is it interesting?

- 1 How many bytes per CU available?
- 2 Schedule on CUs with large sL1 cache

Microbenchmark deep-dive: AMD scalar L1 cache sharing

AMD CDNA architecture



Why is it interesting?

- 1 How many bytes per CU available?
- 2 Schedule on CUs with large sL1 cache
- 3 Cross-CU data exchange

Validation

Tested and verified results on 10 GPUs

① NVIDIA Pascal Quadro P6000

② NVIDIA Volta V100

③ NVIDIA Turing T1000
NVIDIA Turing RTX 2080 Ti

④ NVIDIA Ampere A100

⑤ NVIDIA Hopper H100 80 GB
NVIDIA Hopper H100 96 GB

① AMD CDNA Instinct MI-100

② AMD CDNA2 Instinct MI-210

③ AMD CDNA3 Instinct MI-300X VF

Validation

NVIDIA H100

Memory Element	Size	LD lat	R&W BW	cache line size	fetch granularity	Amount per SM/GPU	Physical hardware sharing
L1 cache	<u>238 KiB</u>	38 ns		128 B	<u>32 B</u>	1	RO, TX, L1
L2 cache	<i>50 MB</i>	220 ns	4.4/3.4 TiB/s	128 B	<u>32 B</u>	2	
Texture cache	<u>238 KiB</u>	<u>39 ns</u>		<u>128 B</u>	<u>32 B</u>	1	RO, TX, L1
Readonly cache	<u>238 KiB</u>	<u>35 ns</u>		<u>128 B</u>	<u>32 B</u>	1	RO, TX, L1
Constant L1 cache	2 KiB	21 ns		64 B	<u>64 B</u>	1	no
Constant L1.5 cache	> 64 KiB	105 ns		×	<u>256 B</u>		
Shared Memory	<i>228 KiB</i>	30 ns					
Device Memory	<i>80 GB</i>	843 ns	2.5/2.7 TiB/s				

Bold = new for this GPU; *Italic* = from an API; **Underlined** = new for this vendor; **×** = not available

Validation

AMD MI210

Memory Element	Size	LD lat	R&W BW	cache line size	fetch granularity	Amount per SM/GPU	Physical hardware sharing
(v)L1 cache	16 KiB	125 ns		64 B	<u>64 B</u>	1	
sL1 cache	15.5 KiB	50 ns		<u>64 B</u>	<u>64 B</u>		<u>CU id</u>
L2 cache	<i>8 MB</i>	<u>310 ns</u>	4.2/2.4 TiB/s	128 B	<u>64 B</u>	1	
Local Data Store	<i>64 KiB</i>	55 ns					
Device Memory	64 GiB	<u>738 ns</u>	1.0/0.9 TiB/s				

Bold = new for this GPU; *Italic* = from an API; **Underlined** = new for this vendor;

Use-case: Dynamic GPU Topologies with sys-sage

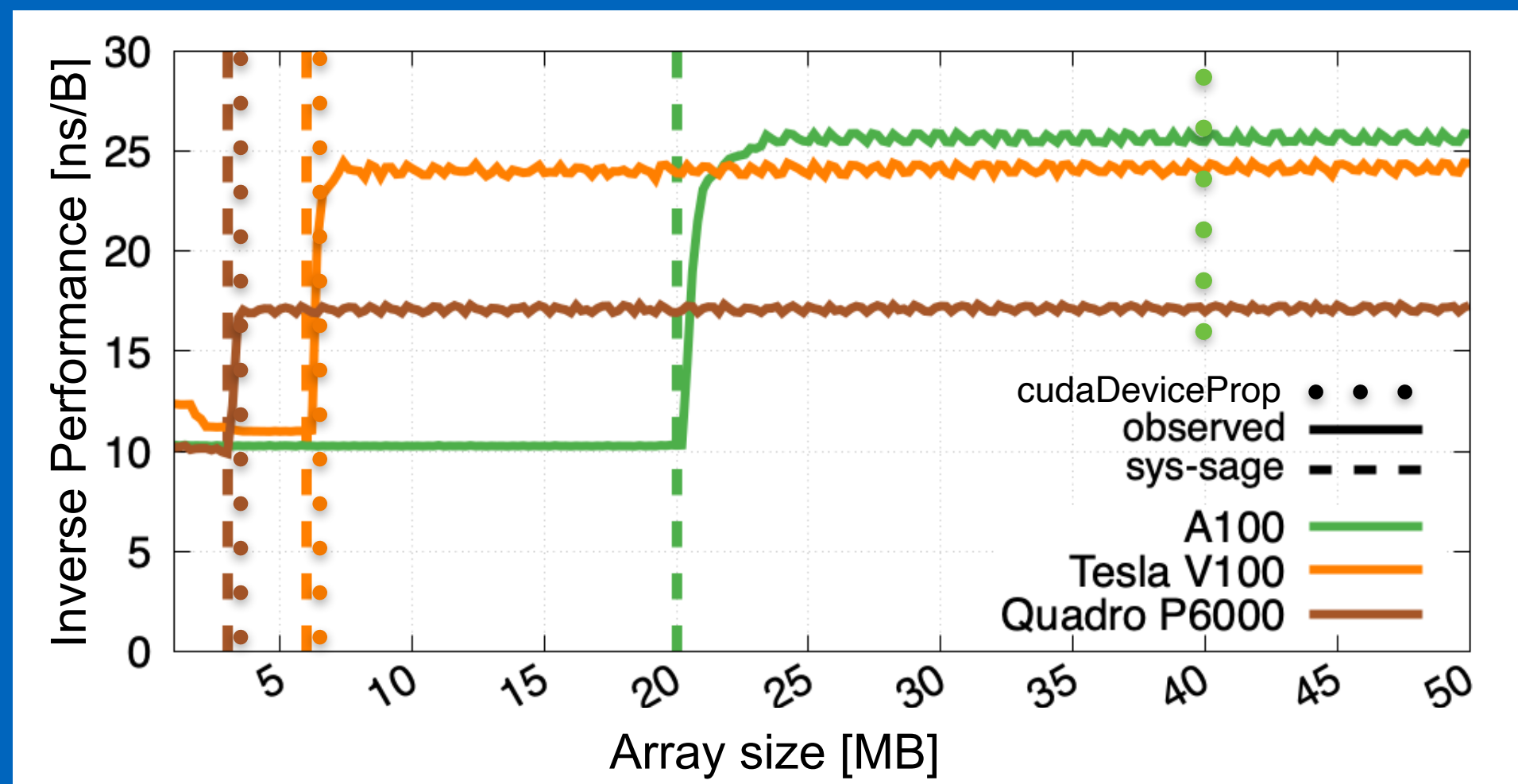
Use-case: Dynamic GPU Topologies with sys-sage

- 
- sys-sage integrates static and dynamic GPU topology information
 - MT4G static topology report
 - NVIDIA `nvm1` API to read current MIG settings

Use-case: Dynamic GPU Topologies with sys-sage

sys-sage integrates static and dynamic GPU topology information

- MT4G static topology report
- NVIDIA `nvm1` API to read current MIG settings

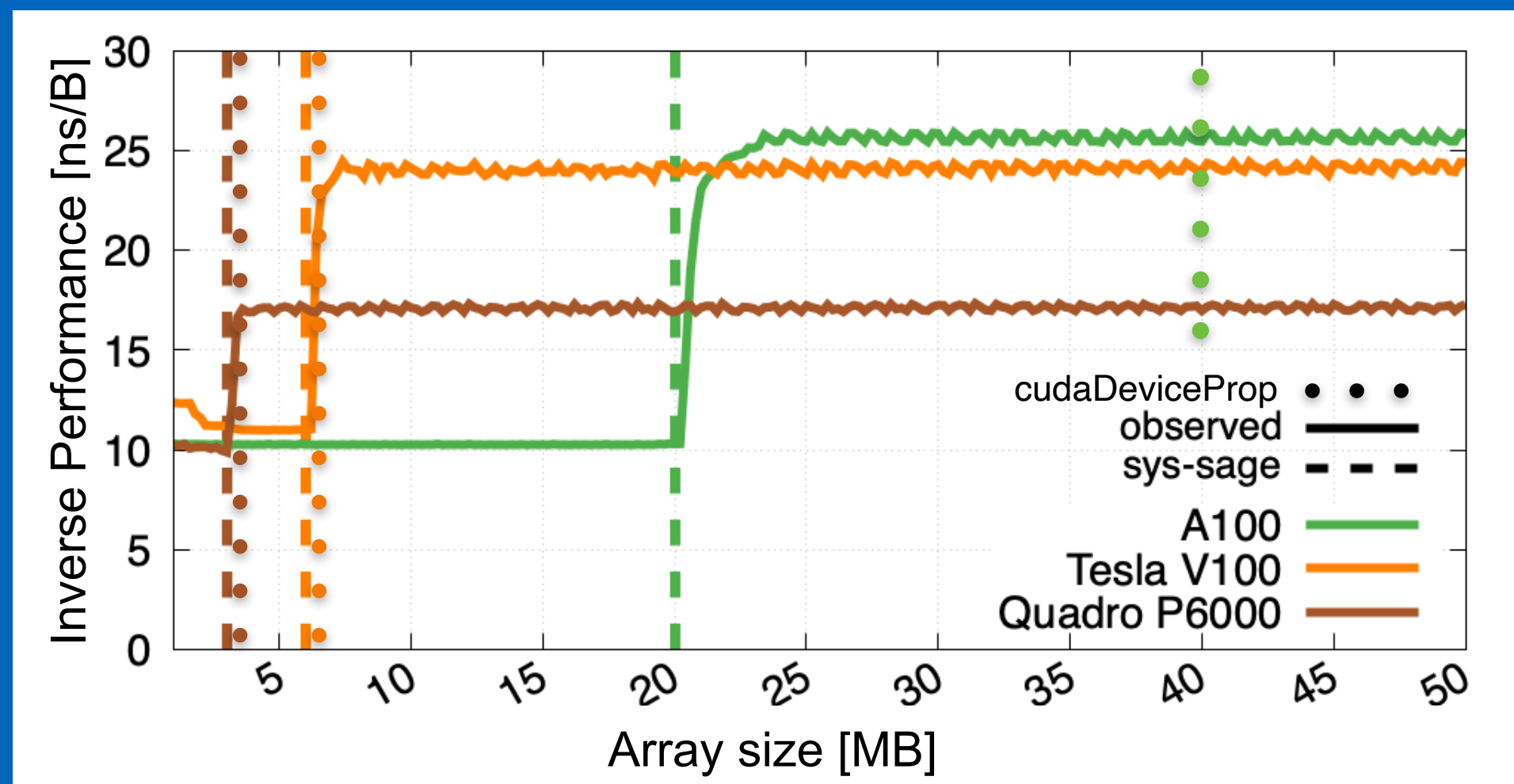


source: Vanecek et al., sys-sage: A Unified Representation of Dynamic Topologies & Attributes

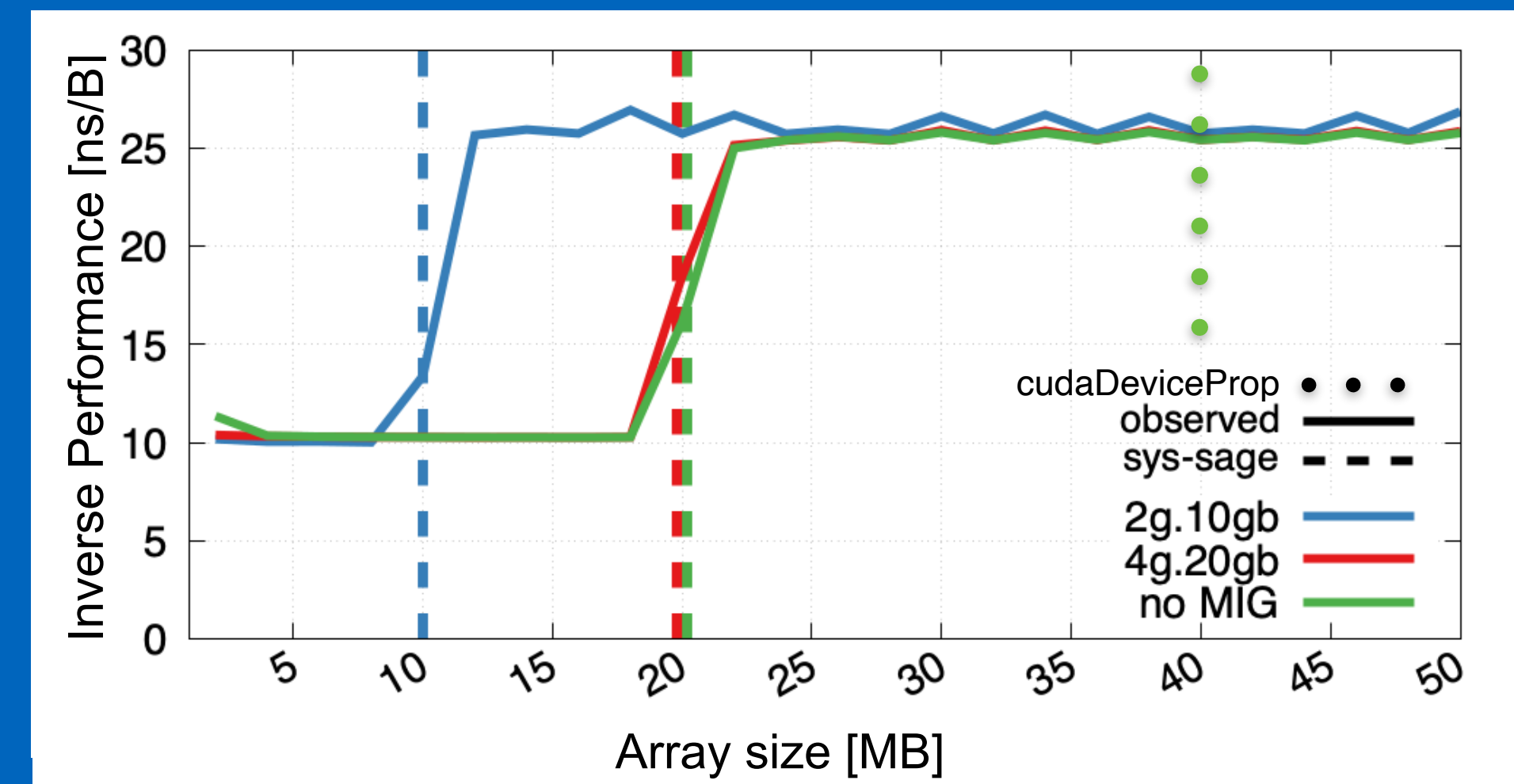
Use-case: Dynamic GPU Topologies with sys-sage

sys-sage integrates static and dynamic GPU topology information

- MT4G static topology report
- NVIDIA `nvm1` API to read current MIG settings



source: Vanecek et al., sys-sage: A Unified Representation of Dynamic Topologies & Attributes



source: Vanecek et al., sys-sage: A Unified Representation of Dynamic Topologies & Attributes

MT4G: GPU Topology Information for NVIDIA and AMD GPUs

MT4G

- Open source
- Works with recent NVIDIA and AMD(CDNA) GPUs
- Provides rich topological information and attributes
- Microbenchmarking at the core
- Automatic evaluation (K-S test)
- Verified on 10 GPUs and 3 use-cases
- Machine+Human readable JSON output

Provides information on

- General GPU information
- Compute Topologies
- Cache and Memory subsystem
 - Size, cache line size, HW layout, ...
 - Load latency, R&W bandwidth

Try MT4G out and reach out to us!



<https://github.com/caps-tum/mt4g>



stepan.vanecek@tum.de



Acknowledgements



 **SEANERGY**
ENERGY EFFICIENT EXASCALE

PDExa

 **PLASMA**
PEPSC

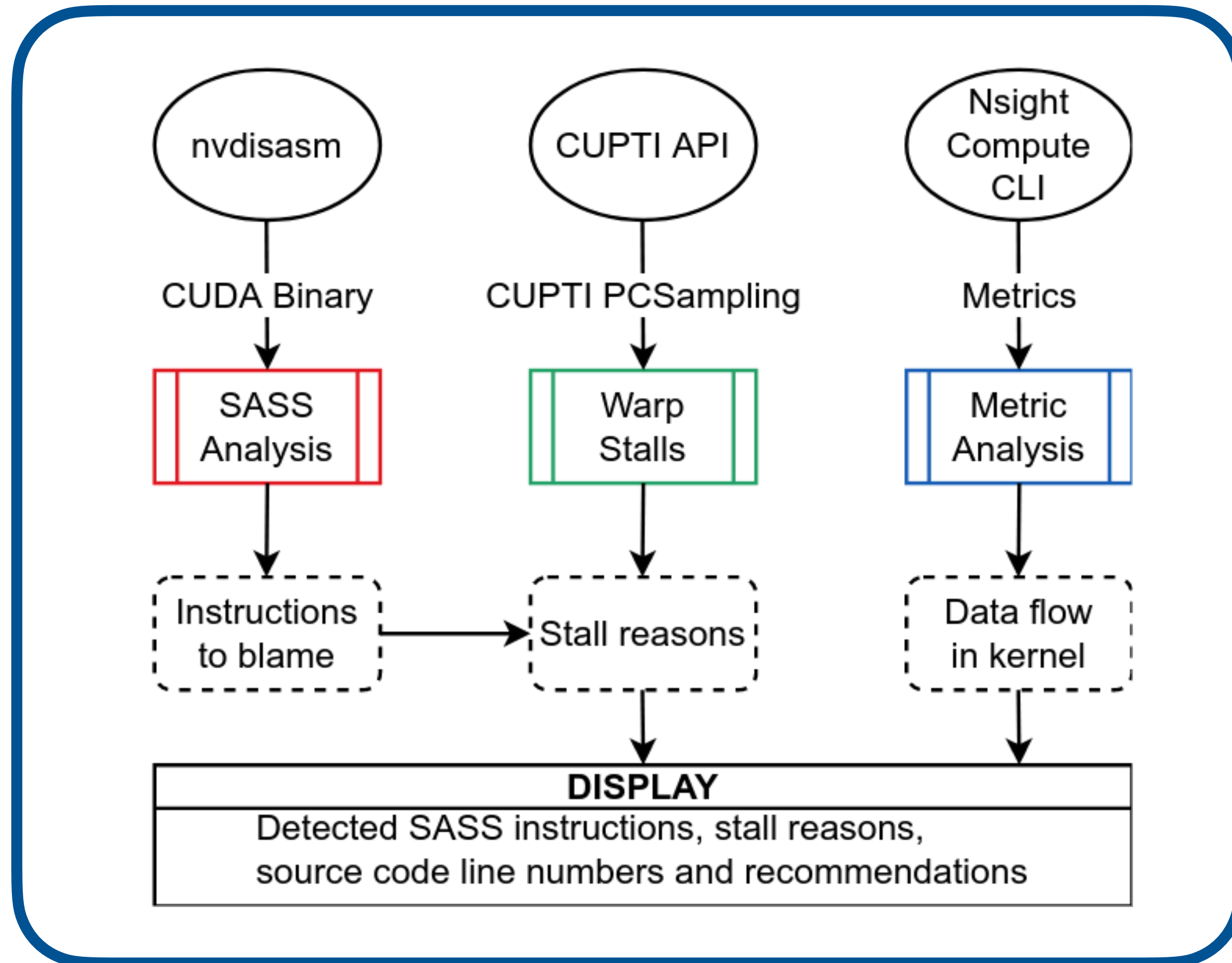
 **DEEP-SEA**



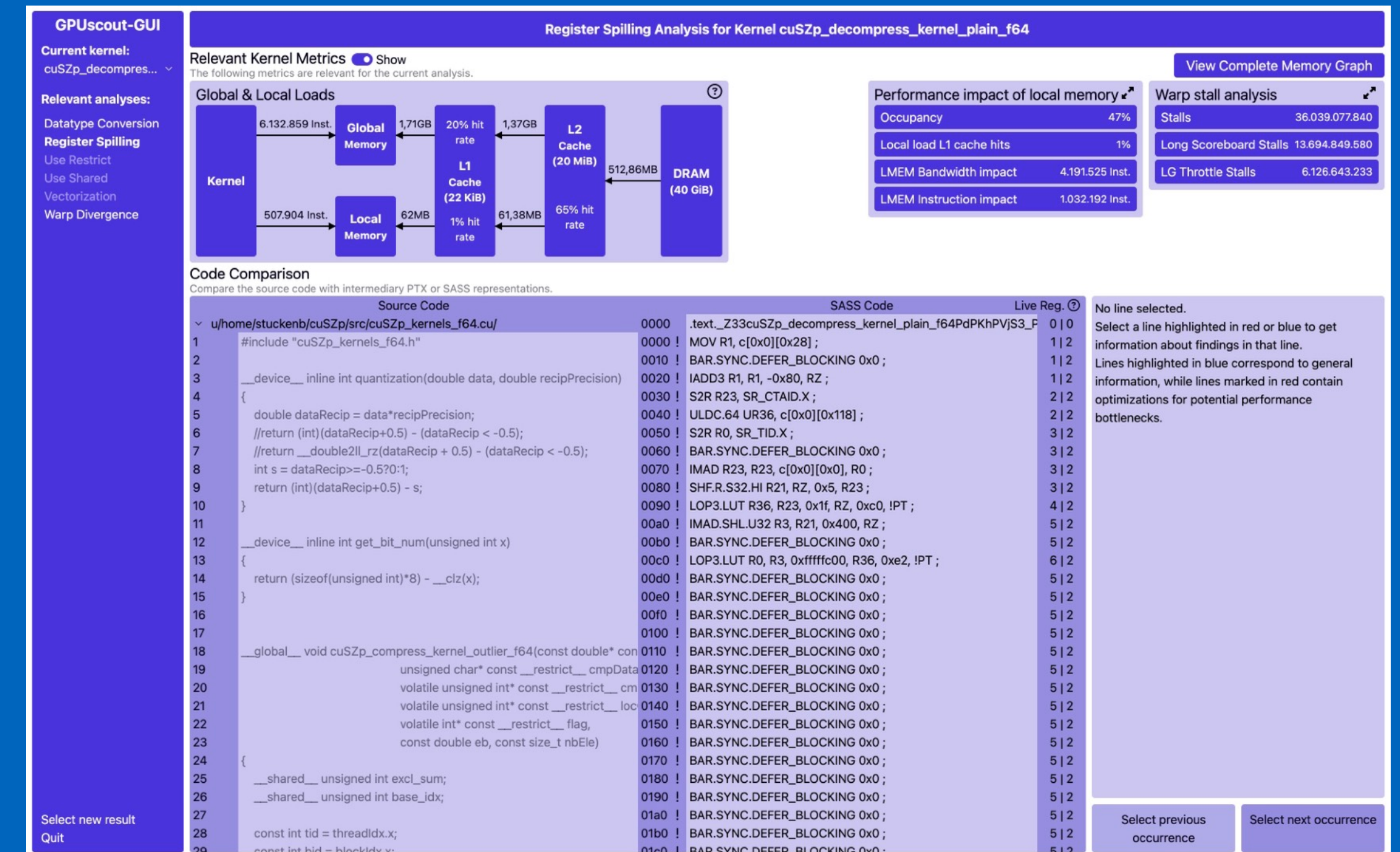
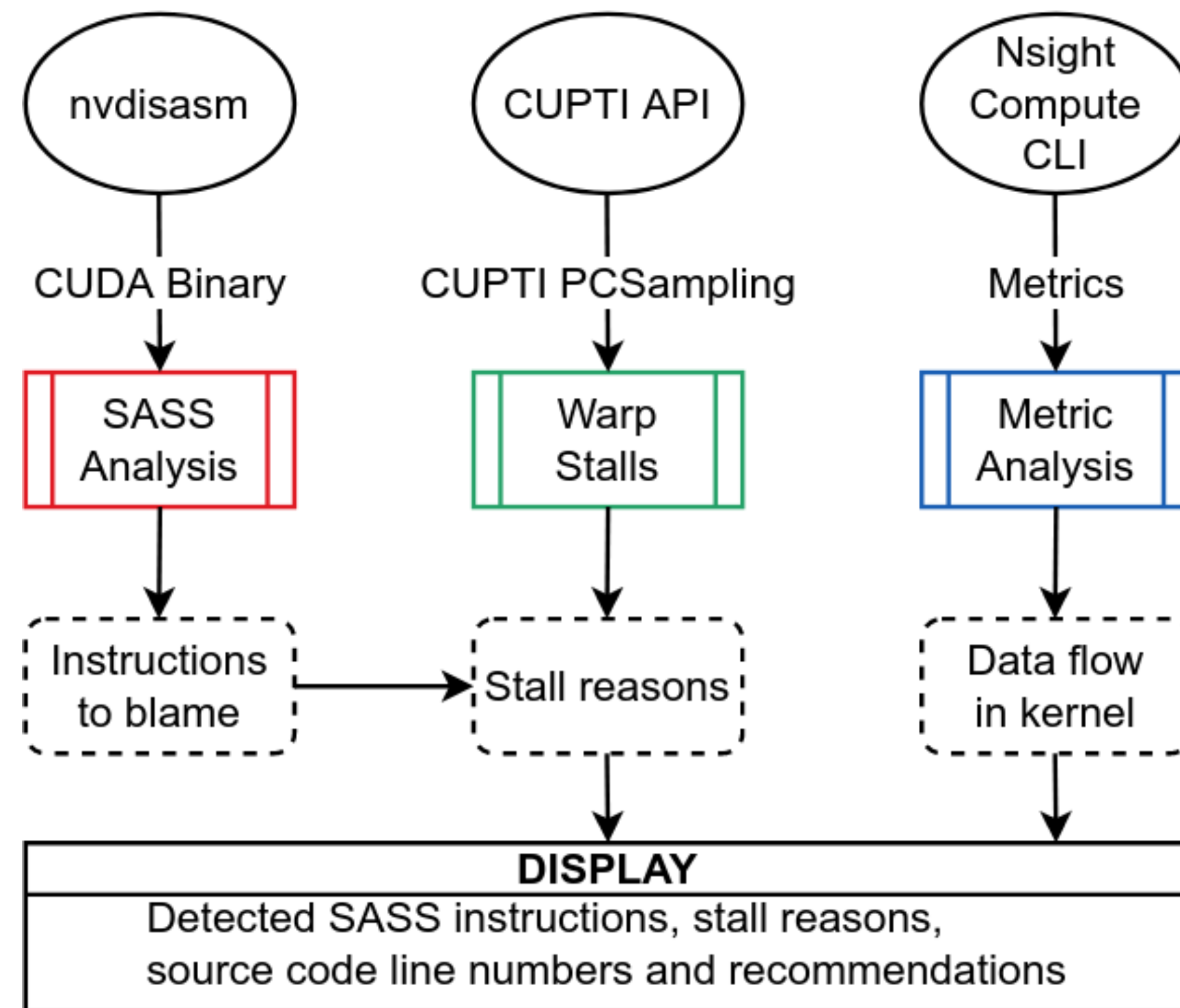
SPONSORED BY THE
Federal Ministry
of Education
and Research

Use-case: Adding HW context to GPUscout

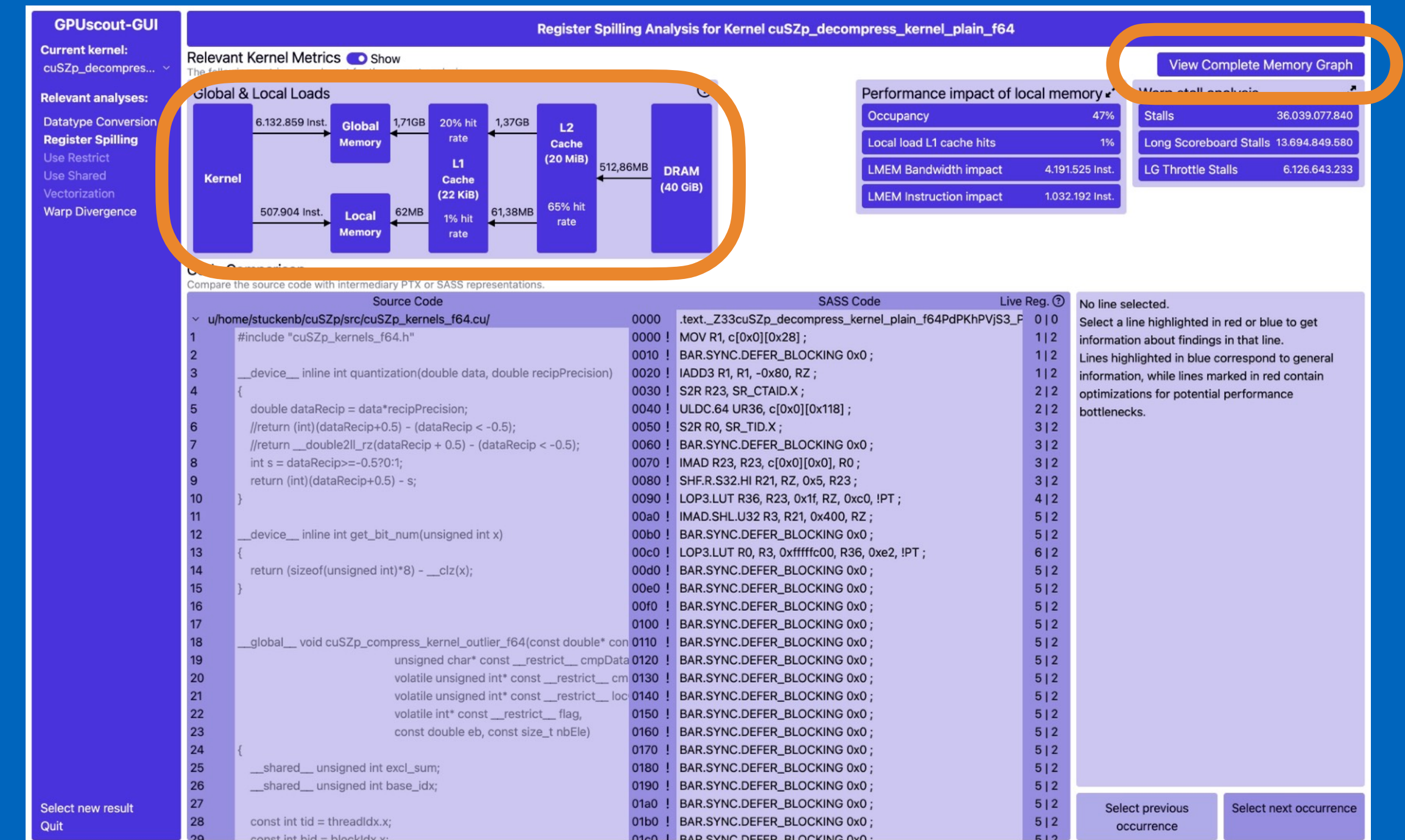
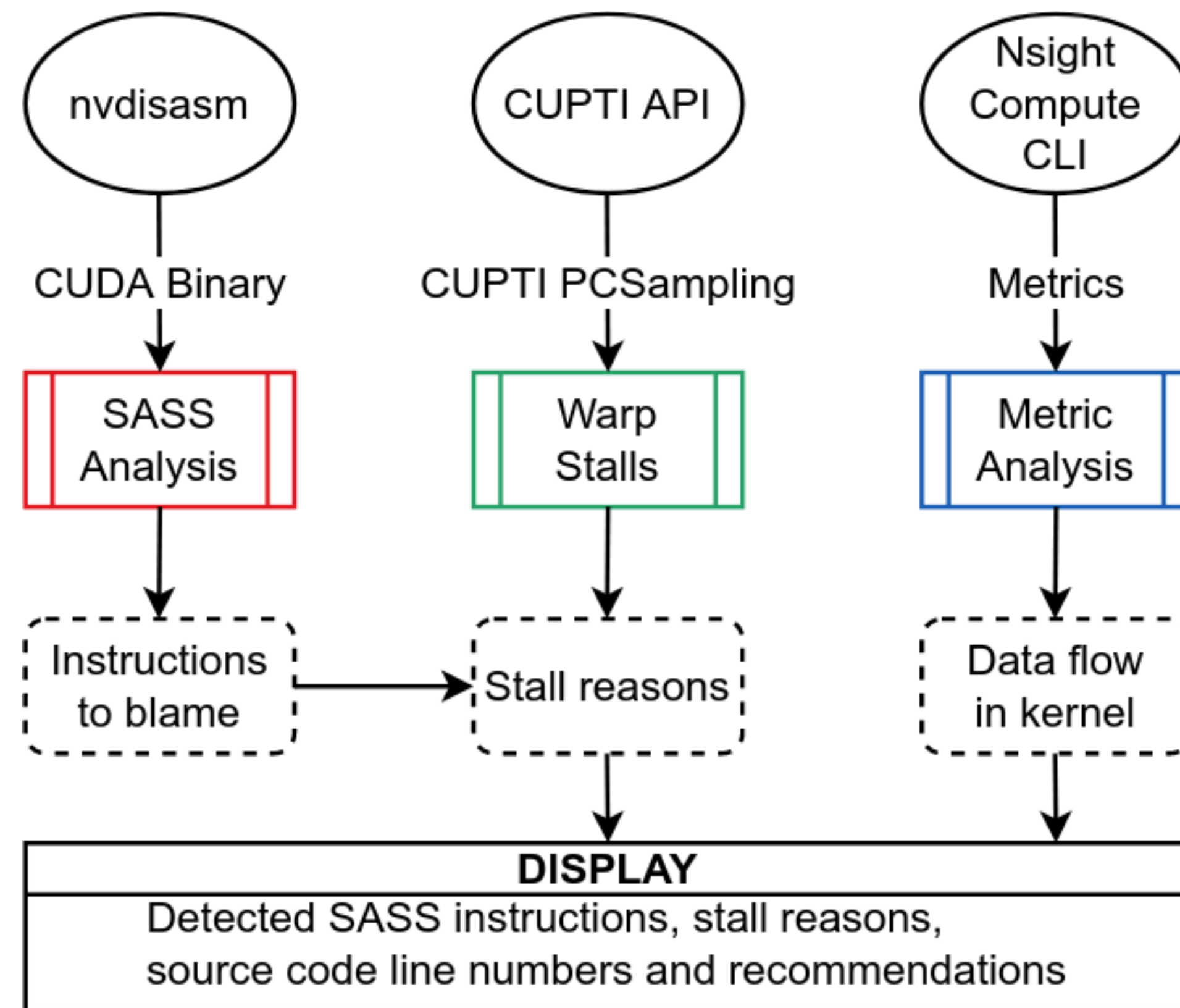
Use-case: Adding HW context to GPUscout



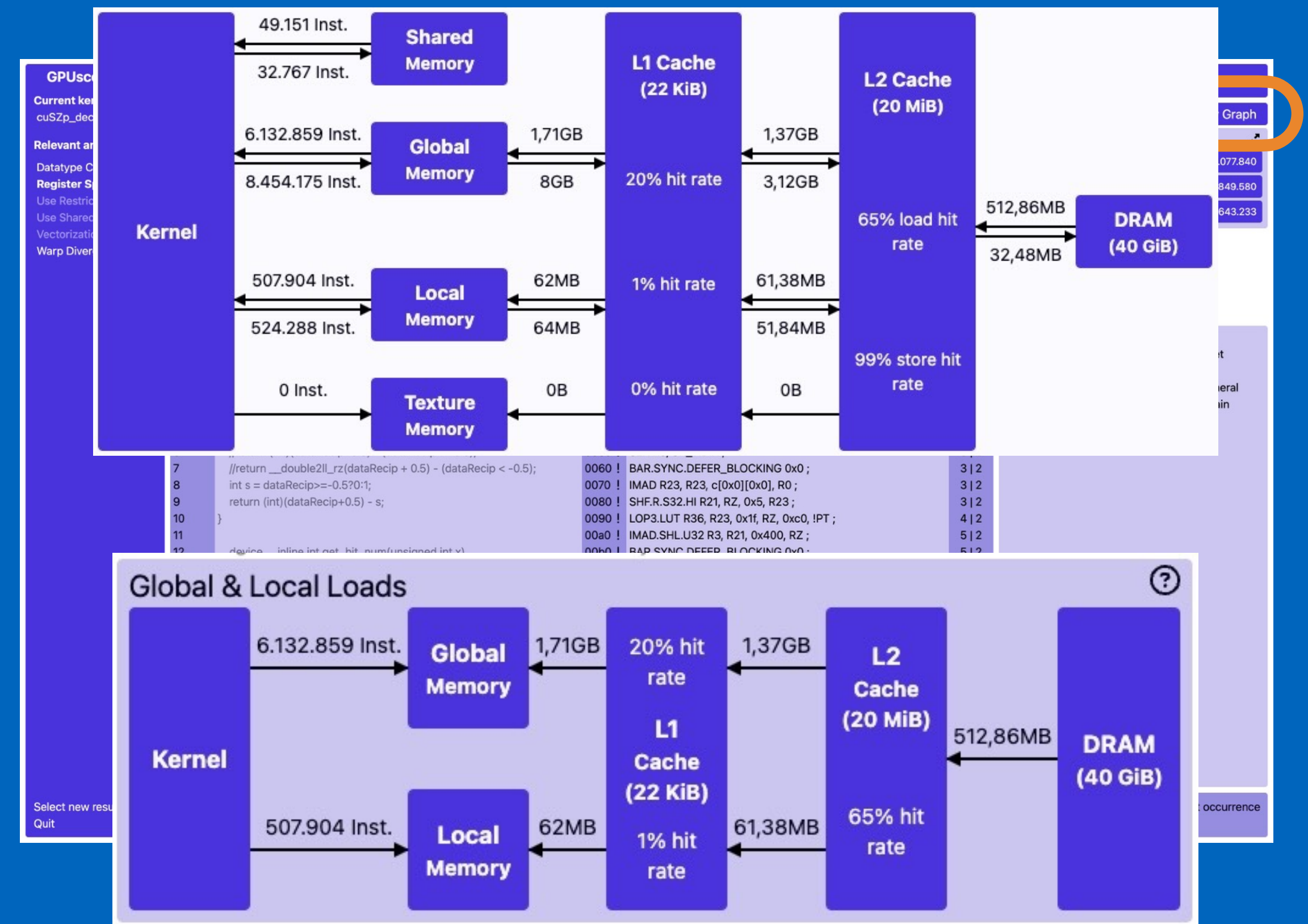
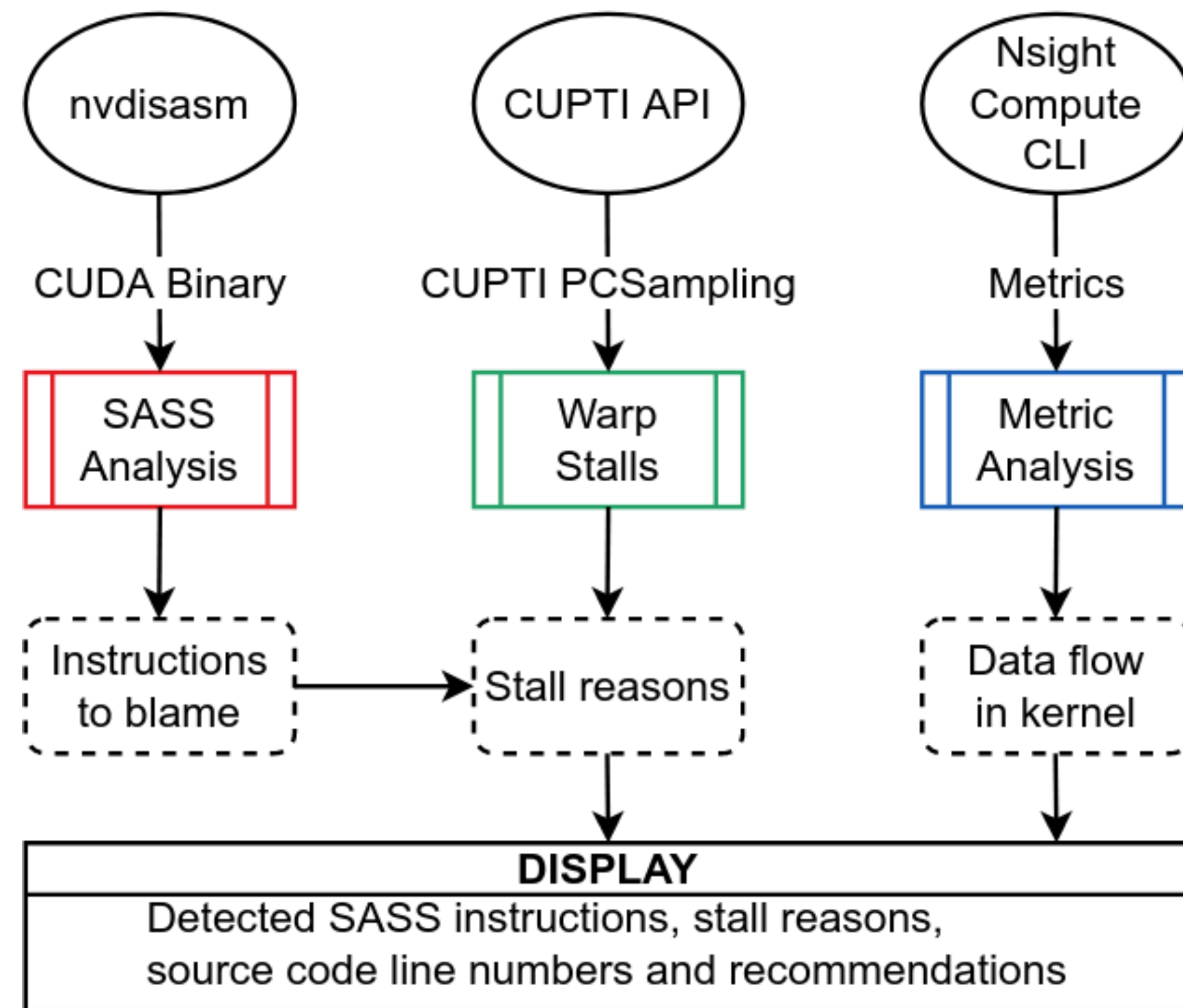
Use-case: Adding HW context to GPUscout



Use-case: Adding HW context to GPUscout



Use-case: Adding HW context to GPUscout



Use-case: Adding HW context to GPUscout

