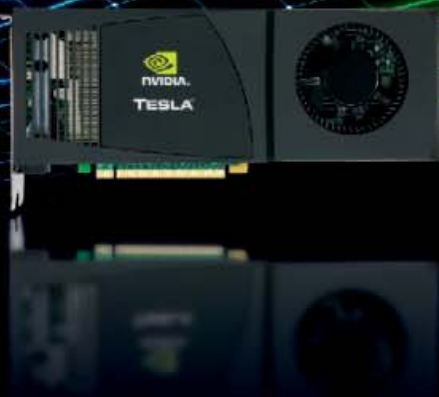




SIGGRAPH Asia 2009 GPU Computing master Class

Future Direction in GPU Computing

Toru Baji / NVIDIA Solution Architect

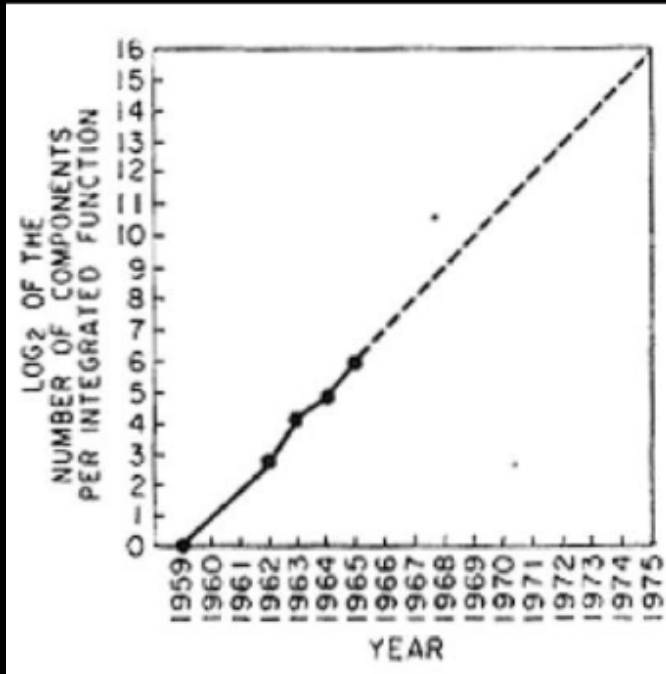
16th Dec., 2009

Content

- **Why GPU Computing?**
 - Single-threaded CPU performance is no longer scaling
 - CUDA can exploit “ Performance = Parallelism”
- **Evolution of NVIDIA GPU**
- **Key Milestone in GPU Computing ‘Fermi’**
 - Processor / Memory Architecture
 - Enhanced Thread Management
 - CUDA3.0
 - Integrated Development Environment ‘Nexus’
- **An NVIDIA ExaScale Machine in 2017**
- **Conclusion**

Single-threaded CPU performance is no longer scaling (1)

Moore's Law



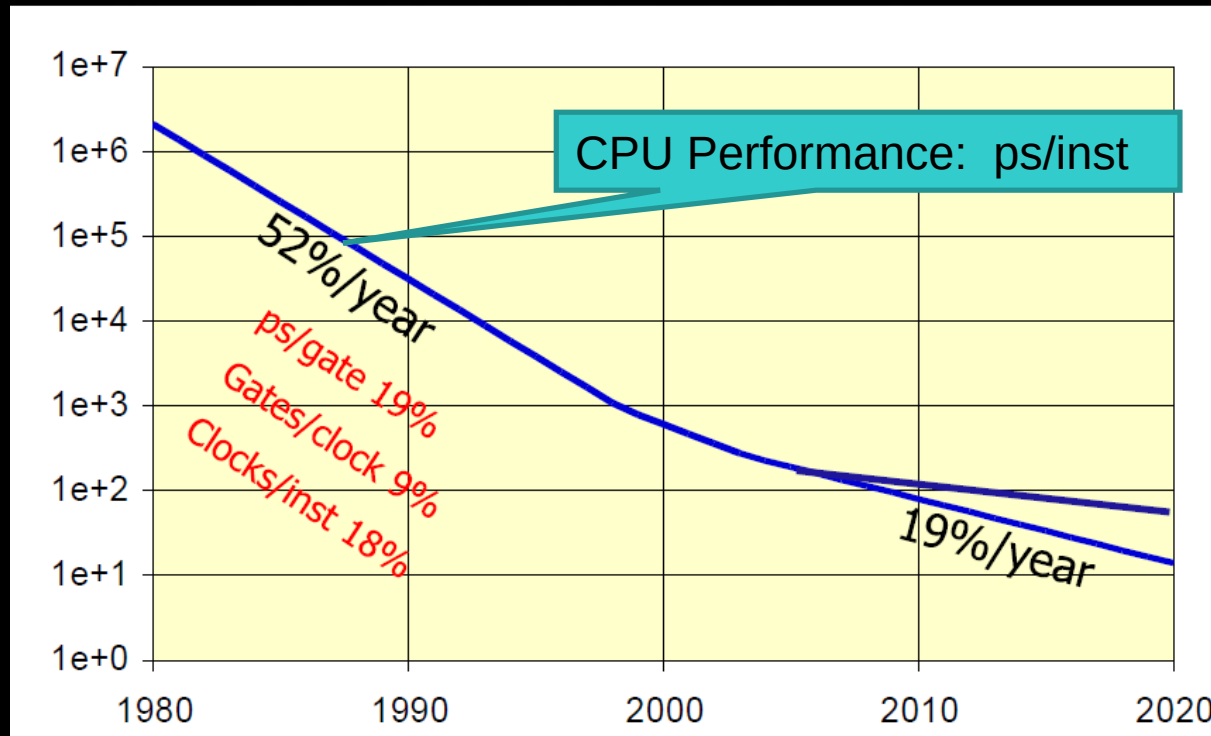
Moore, Electronics 38(8) April, 1965

- In 1965, Gordon Moore predicted the number of Tr will grow at the rate of:
x 2 / year
later revised to
x 2 / 18-months
- No prediction of CPU performance

Single-threaded CPU performance is no longer scaling (2)

The End of ILP (Instruction Level Parallelism) Scaling

- The increase in Tr counts was used to increasing the parallelism of the CPU instruction execution (e.g. superscalar, pipeline) and other speed-ups



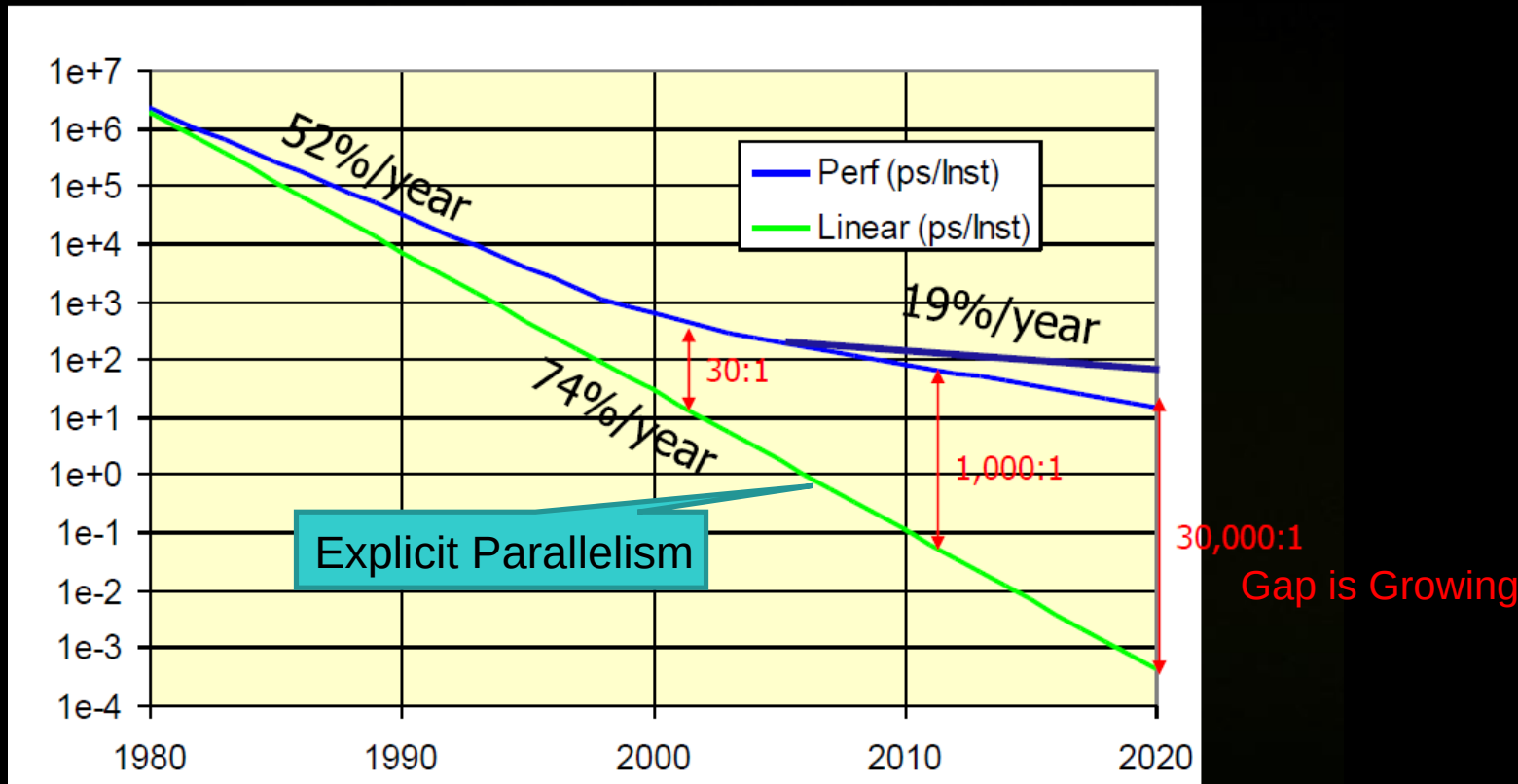
Performance is measured as the time required to execute one instruction (ps/inst)

Bill Dally et al., the Last Classical Computer, ISAT Study, 2001

Single-threaded CPU performance is no longer scaling (3)

Explicit Parallelism is Now Attractive

- The increase in Tr counts is used to increasing the number of processors



Bill Dally et al., the Last Classical Computer, ISAT Study, 2001

CUDA can exploit Performance = Parallelism



Now it is clear that “ **Performance = Parallelism**”

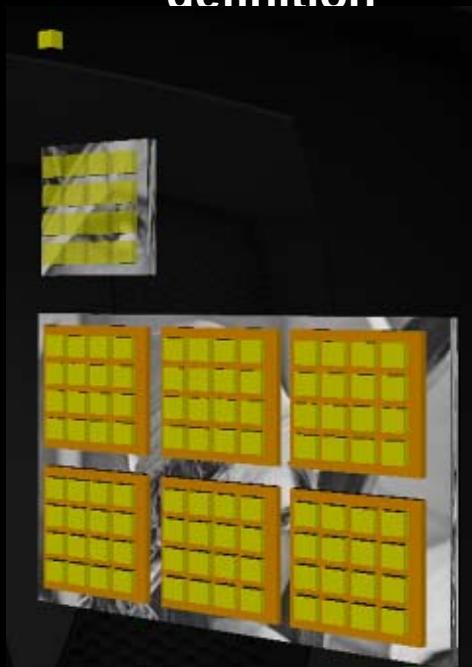
- Then what is required to exploit this Parallelism?
 - Many efficient processors
 - A programming system that abstract the parallelism
- “**CUDA** (Compute Unified Device Architecture)”
 - GPU has hundreds of full-featured processors
 - Multi-threading architecture use the best of those massive parallel cores
 - CUDA starts from extension to the familiar C-language, and now expanded to Fortran, OpenCL and direct Compute
 - CUDA can isolate the programmer from the details of parallel programming

CUDA GPU Computing Visual Demo

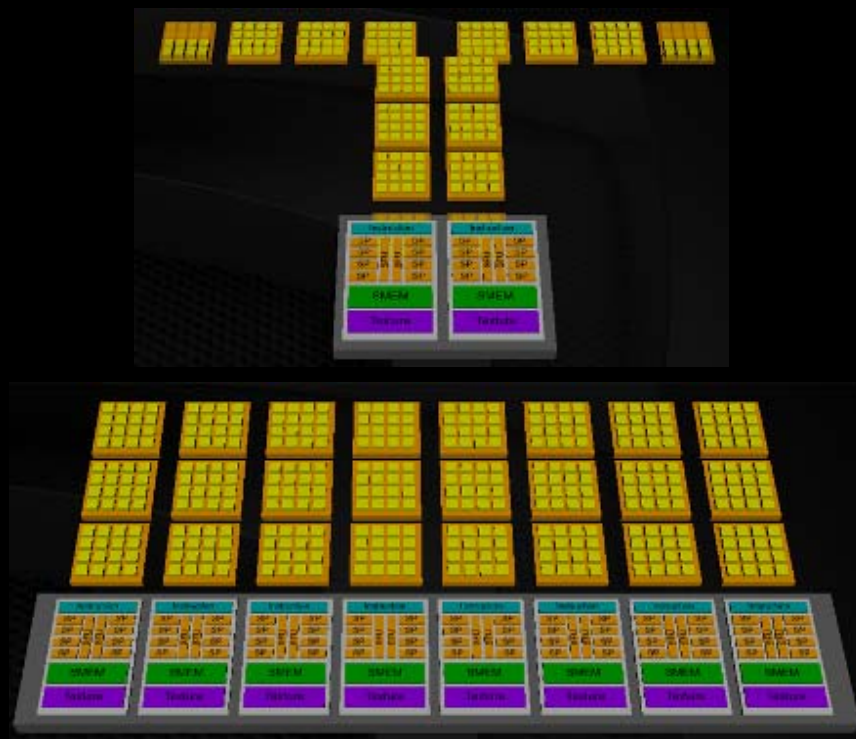
separate animation tool used



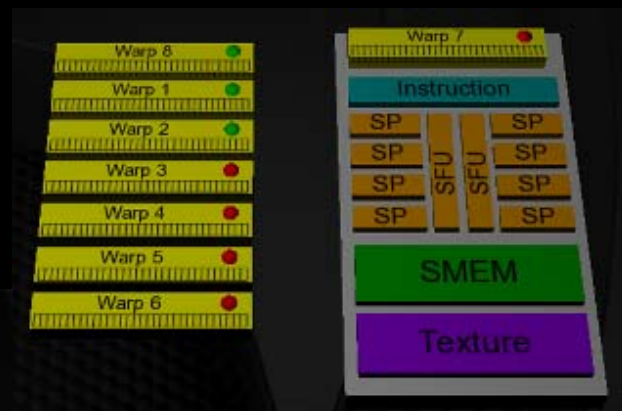
Hierarchical thread definition



Scalability to n-SMs



Hide memory Latency



Ease of Programming and Performance Increase

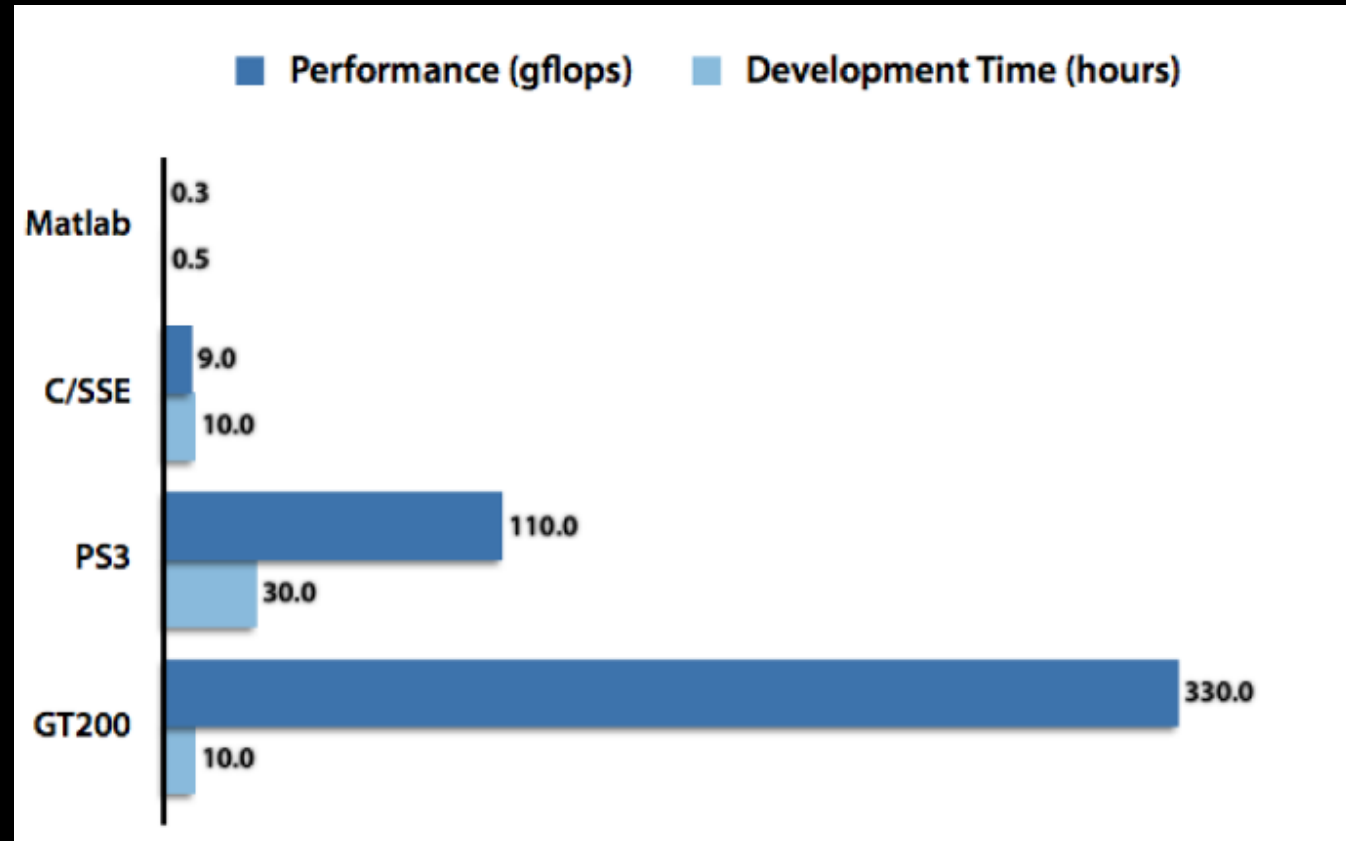
- Ease of programming represented as “Development time”
- Performance achievement measured in GFLOPS

High-level algorithm
Development tool on CPU

CPU & SSE-SIMD

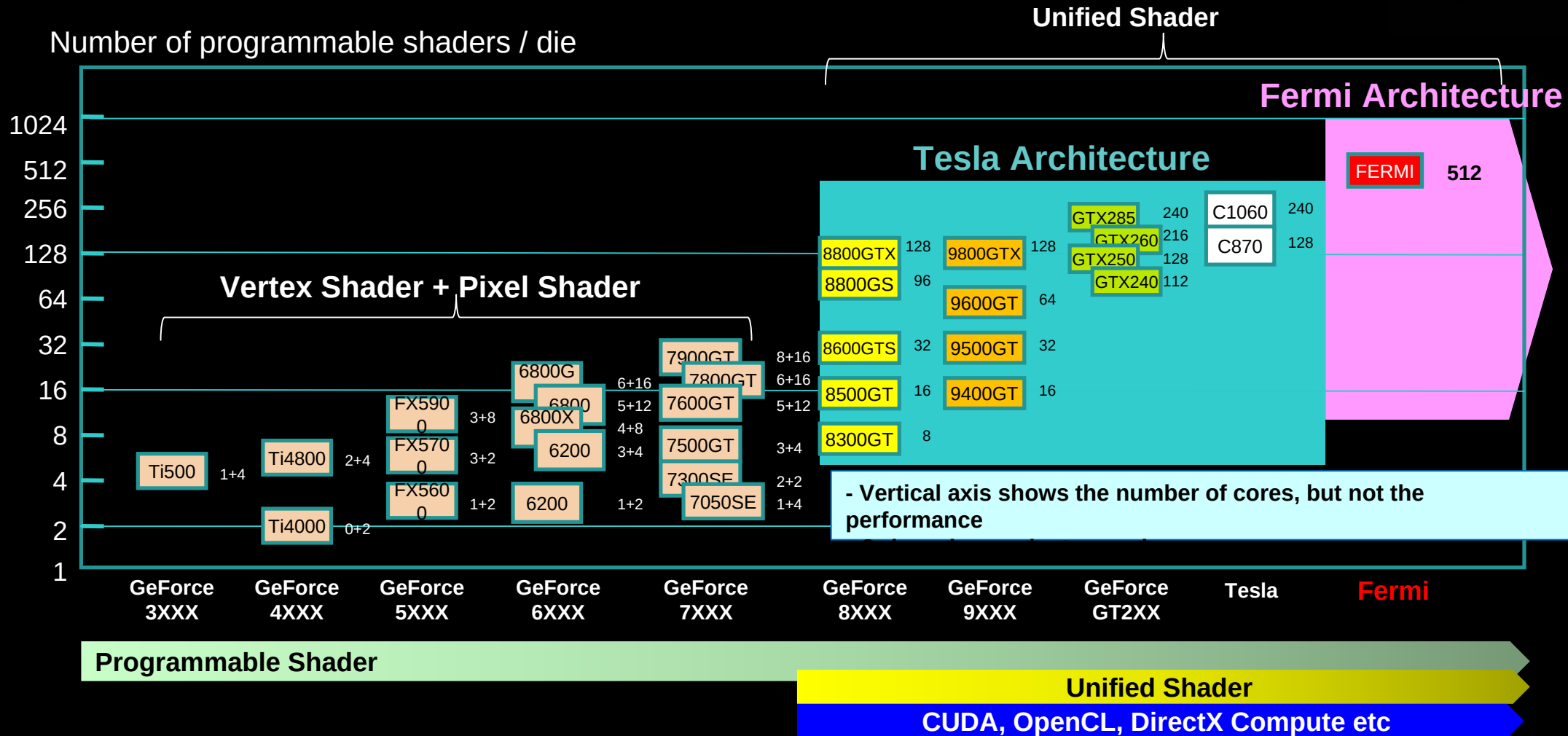
Cell Processor

GPU with CUDA



Source: Nicolas Pinto, MIT

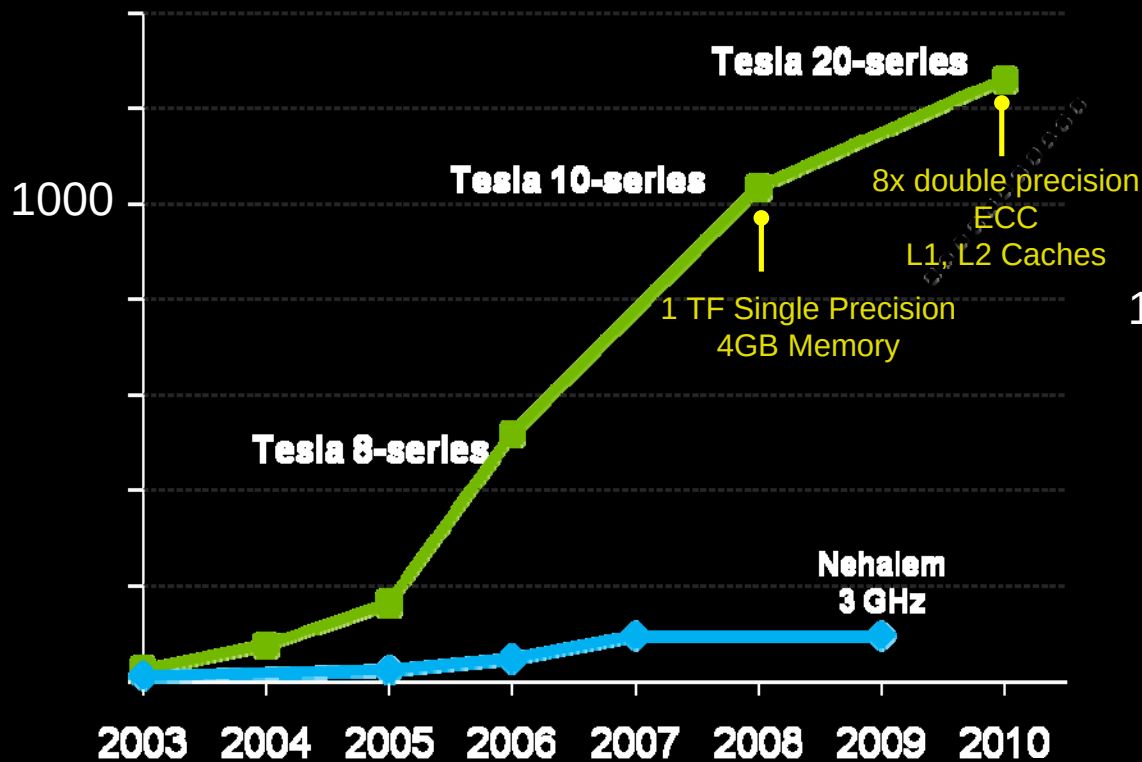
Evolution of NVIDIA GPU - Number of Cores



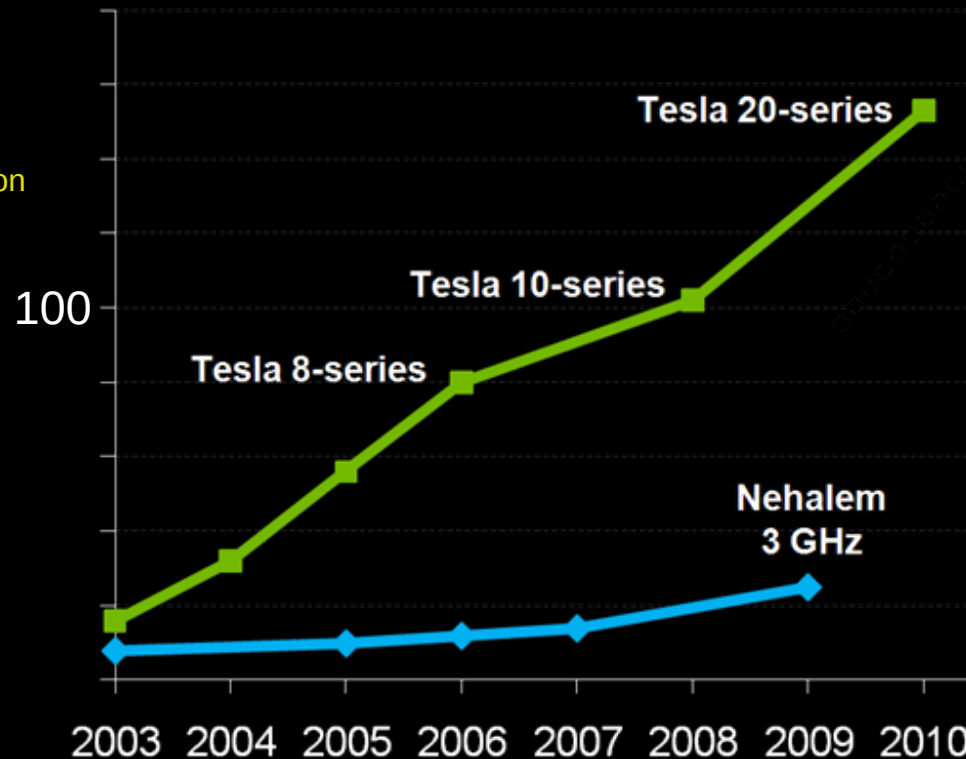
Gap of GPU-CPU Performance, Mem-BW is growing



Peak Single Precision Performance GFlops/sec



Peak Memory Bandwidth GB/sec

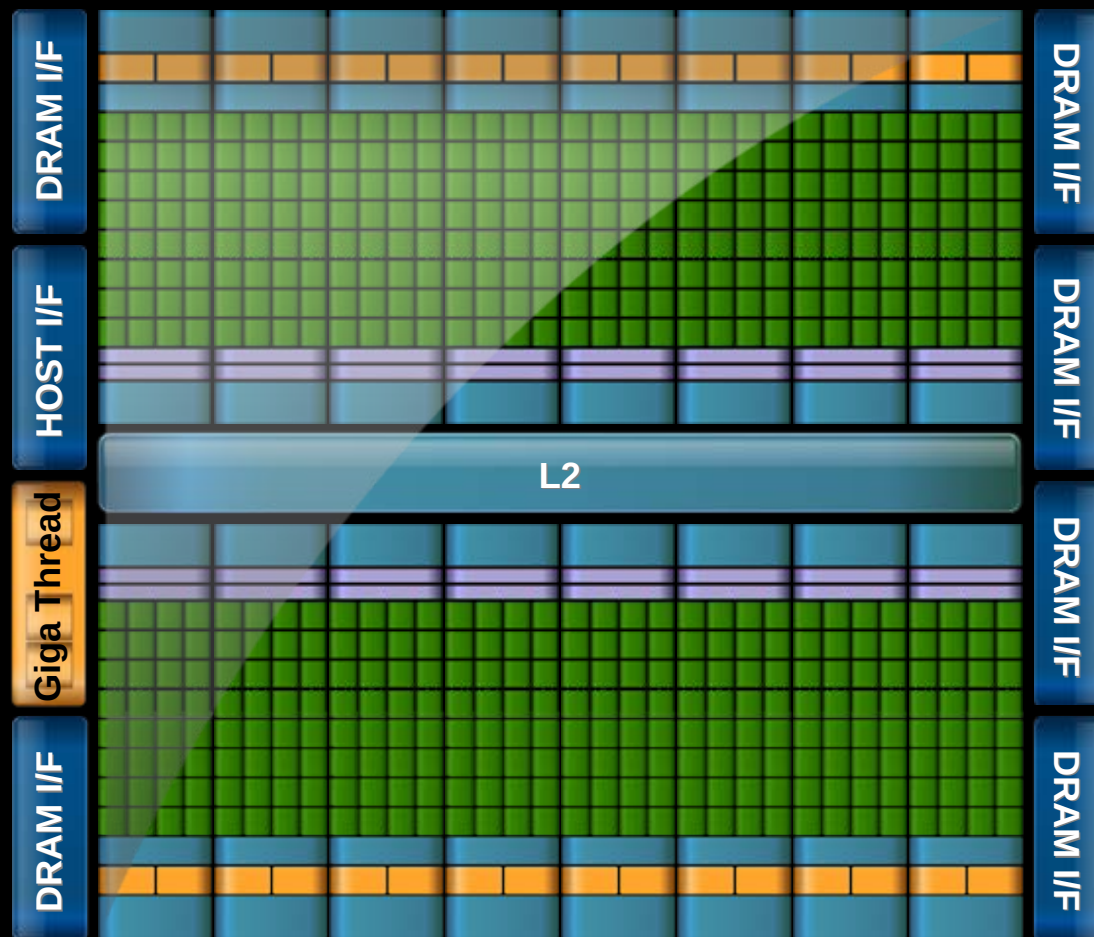


■ NVIDIA GPU
◆ X86 CPU

Key Milestone in GPU Computing 'Fermi'



- Improved peak performance
- Improved throughput efficiency
- Broader applicability
- Full integration within modern software development environment

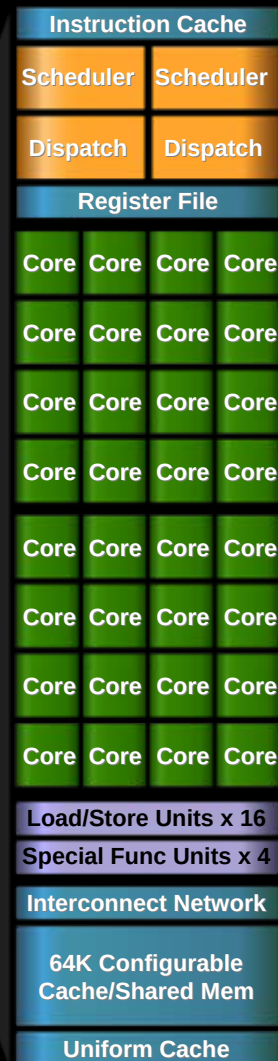


SM Multiprocessor Architecture

- 32 CUDA cores per SM (512 total)
- 8x peak FP64 performance
 - 50% of peak FP32 performance
- Dual Thread Scheduler
- 64 KB of RAM for shared memory and L1 cache (configurable)



	FP32	FP64	INT	SFU	LD/ST
Ops / clk	32	16	32	4	16



IEEE 754-2008 Floating Point



- **IEEE 754-2008 results**

- 64-bit double precision
- 32-bit single precision
- full-speed denormal operands & results
- NaNs, +/- Infinity

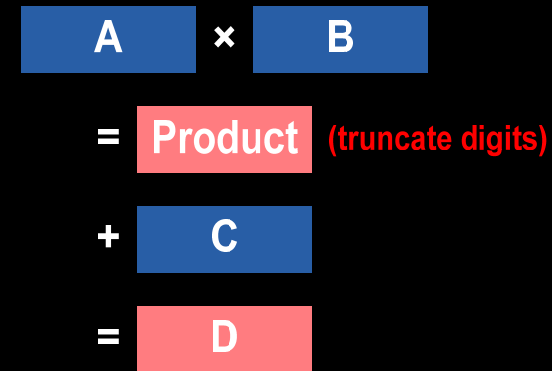
- **IEEE 754-2008 rounding**

- nearest even, zero, +inf, -inf

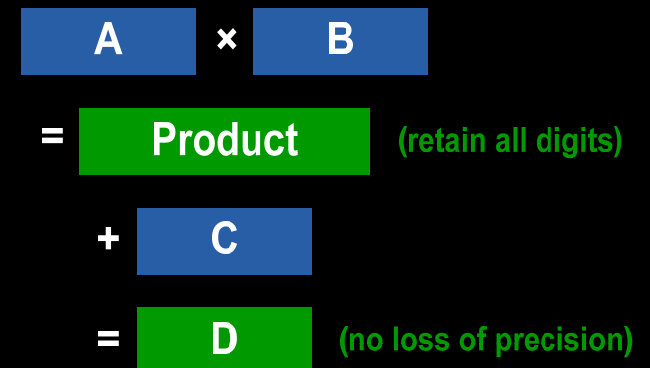
- **IEEE 754-2008 Fused Multiply-Add (FMA)**

- $D = A * B + C$;
- No loss of precision
- IEEE divide & sqrt use FMA

Multiply-Add (MAD): $D = A * B + C$;



Fused Multiply-Add (FMA): $D = A * B + C$;

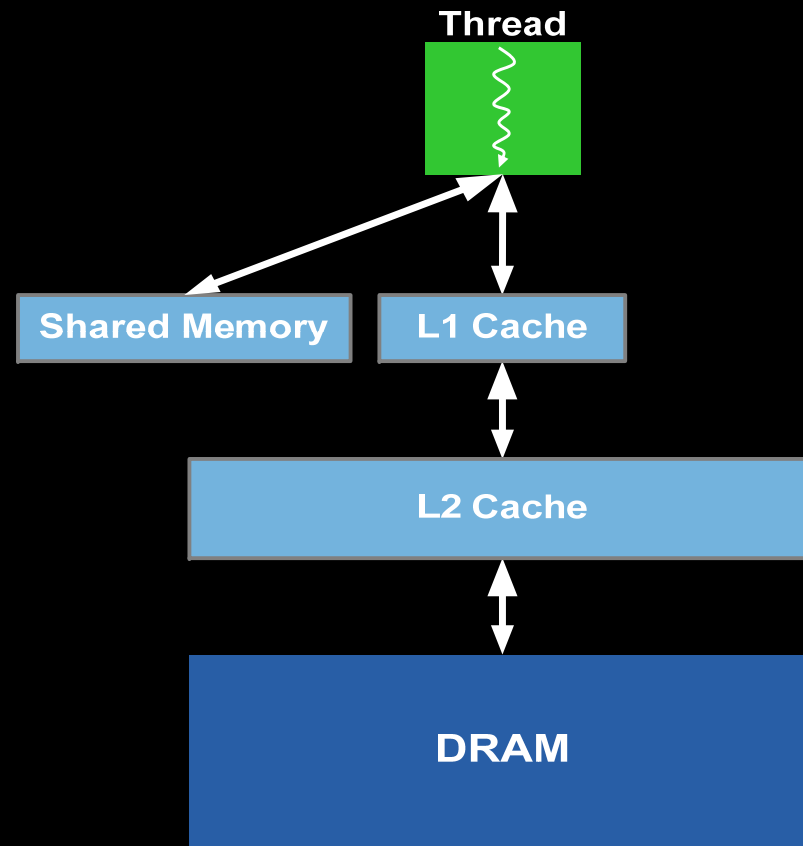


Cached Memory Hierarchy



- **Configurable L1 cache per SM**
 - 16KB L1\$ / 48KB Shared Memory
 - 48KB L1\$ / 16KB Shared Memory
- **Shared 768KB L2 cache**
- **Motivation:**
 - Caching captures locality, amplifies bandwidth
 - Caching more effective than Shared Memory RAM for irregular or unpredictable access
 - Ray tracing, sparse matrix multiply, physics kernels ...
 - Caching helps latency sensitive cases

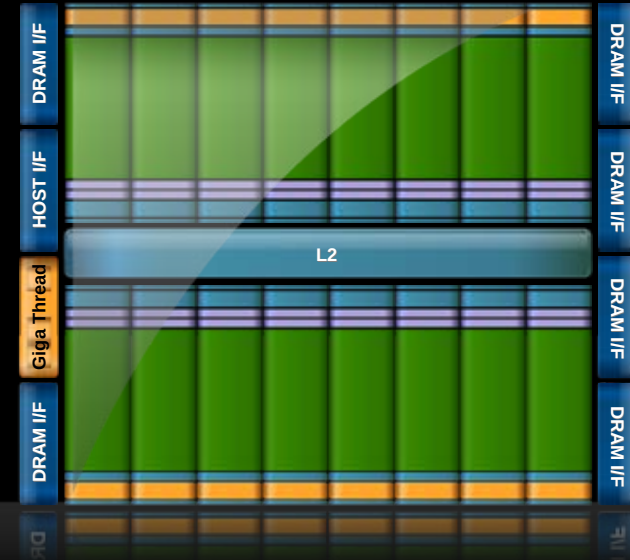
Fermi Memory Hierarchy



Larger, Faster Memory Interface



- **GDDR5 memory interface**
 - 2x improvement in peak speed over GDDR3
- **Up to 1 Terabyte of memory attached to GPU**
 - Operate on large data sets



Unified Load/Store Addressing



Non-unified Address Space



Unified Address Space



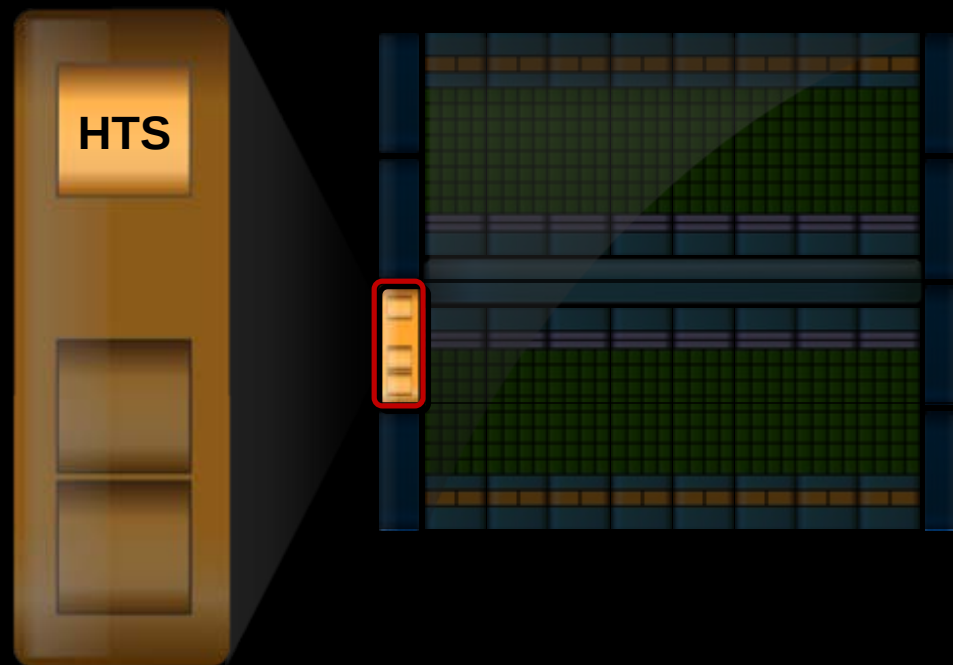
ECC Memory Protection

- **All major internal memories protected by ECC**
 - Register file
 - L1 cache
 - L2 cache
- **External DRAM protected by ECC**
- **ECC is a must have for many computing applications**
 - Clear customer feedback

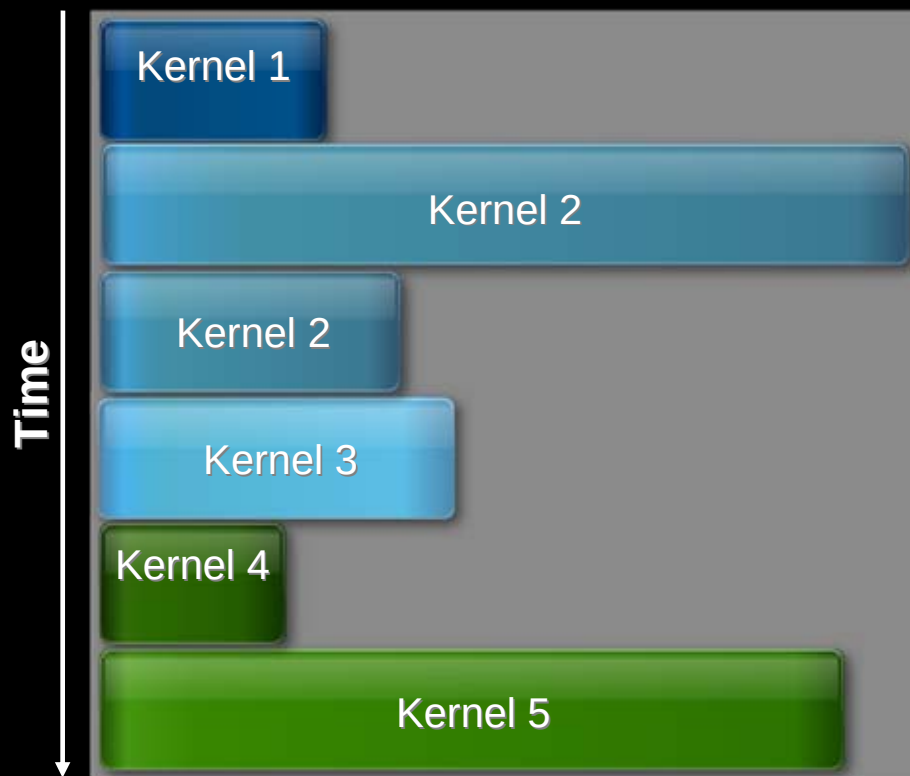
GigaThread™ Hardware Thread Scheduler



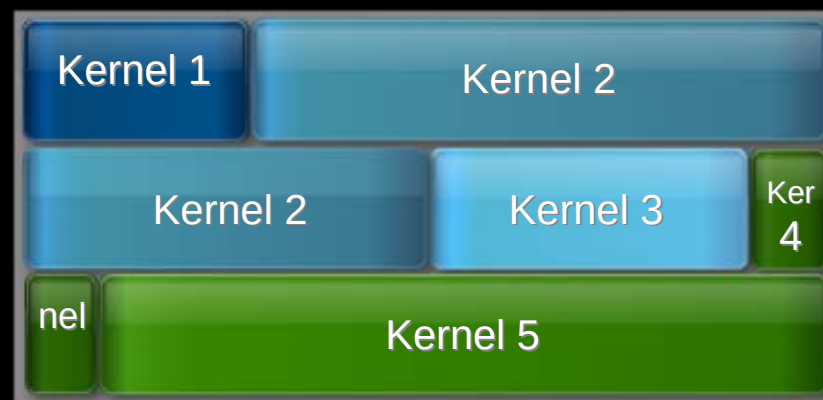
- Hierarchically manages tens of thousands of simultaneously active threads
- 10x faster context switching on Fermi
- Overlapping kernel execution



Overlapping Kernel Execution



Serial Kernel Execution

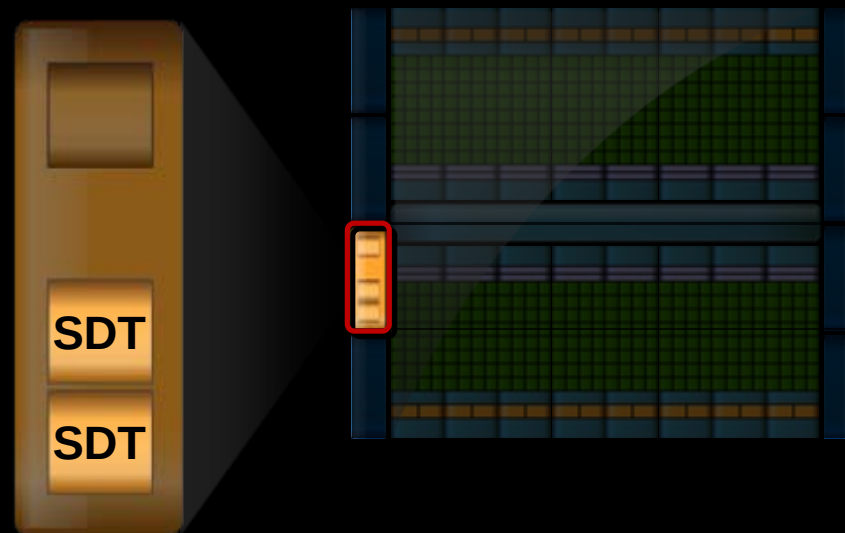


Parallel Kernel Execution

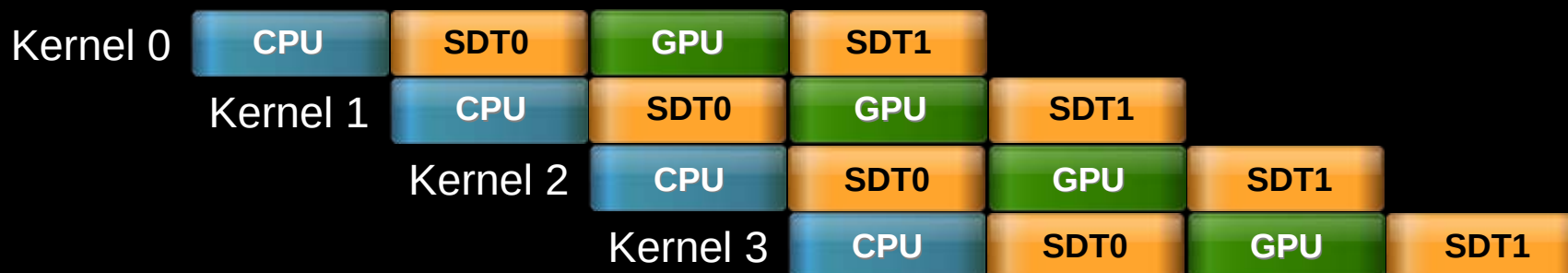
GigaThread Streaming Data Transfer Engine



- Dual DMA engines
- Simultaneous CPU→GPU and GPU→CPU data transfer



- Fully overlapped with CPU/GPU processing



	G80	GT200	Fermi
Transistors	681 million	1.4 billion	3.0 billion
CUDA Cores	128	240	512
Double Precision Floating Point	-	30 FMA ops/clock	256 FMA ops/clock
Single Precision Floating Point	128 MAD ops/clock	240 MAD ops/clock	512 FMA ops/clock
Special Function Units (per SM)	2	2	4
Warp schedulers (per SM)	1	1	2
Shared Memory (per SM)	16 KB	16 KB	Configurable 48/16 KB
L1 Cache (per SM)	-	-	Configurable 16/48 KB
L2 Cache	-	-	768 KB
ECC Memory Support	-	-	Yes
Concurrent Kernels	-	-	Up to 16
Load/Store Address Width	32-bit	32-bit	64-bit

CUDA Parallel Computing Architecture



GPU Computing Applications

C++

C

OpenCL™

Direct
Compute

Fortran

Java and
Python



NVIDIA GPU

with the CUDA Parallel Computing Architecture

CUDA C 3.0

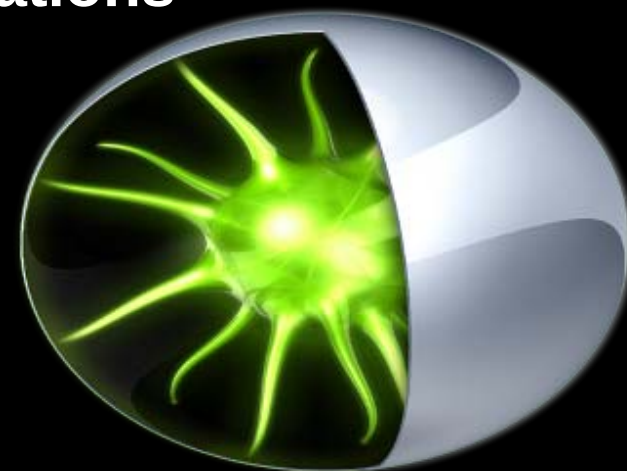


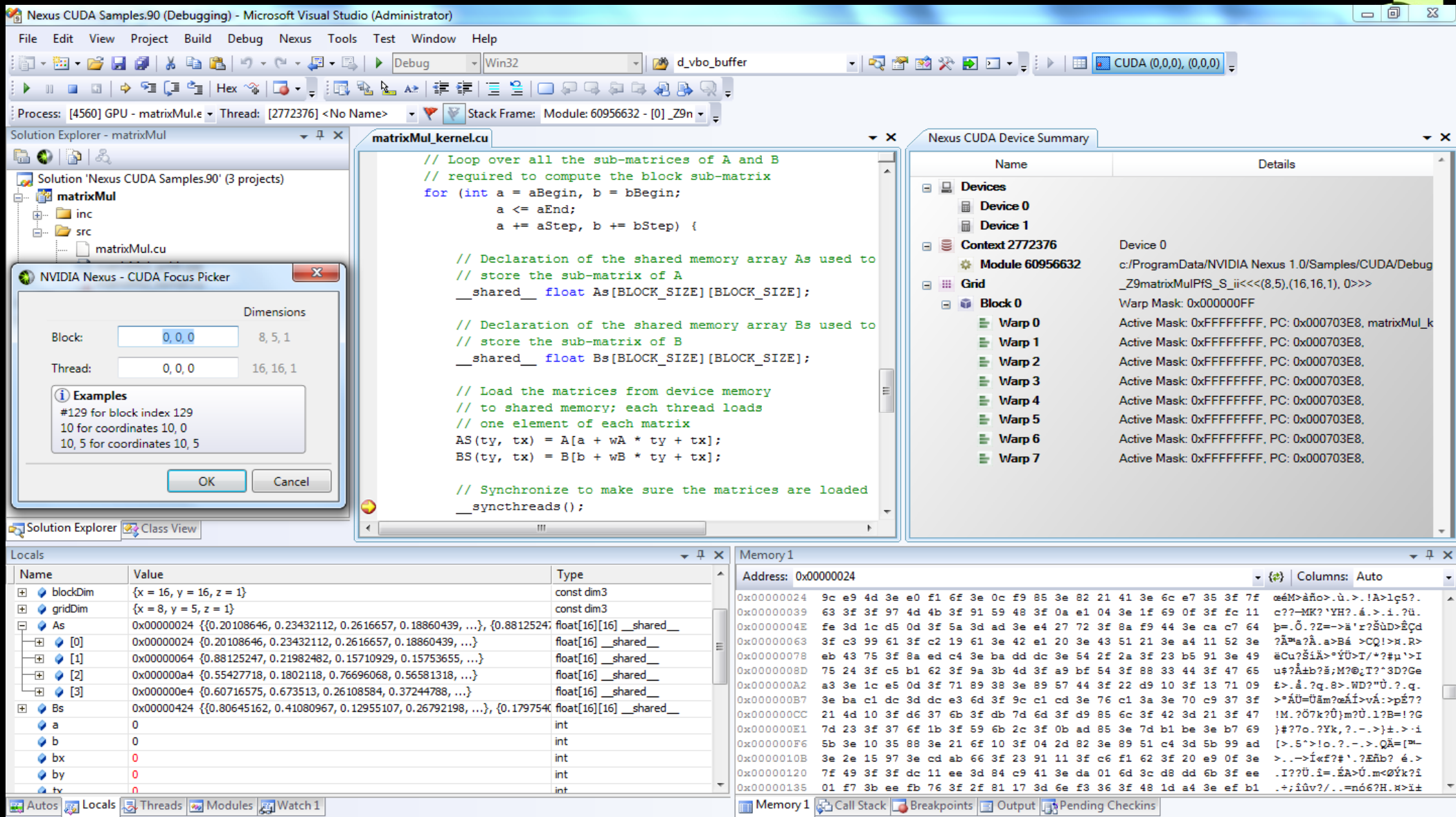
- **Unified addressing for C and C++ pointers**
 - Global, shared, local addresses
 - Enables 3rd party GPU callable libraries, dynamic linking
 - One 40-bit address space for load/store instructions
-
- **Compiling for native 64-bit addressing**
-
- **IEEE 754-2008 single & double precision**
 - C99 math.h support

NVIDIA Integrated Development Environment Code-named 'Nexus'



- Industry's 1st IDE for **massively parallel** applications
- Accelerate development of CPU + GPU **co-processing** applications
- Complete **Visual Studio-integrated** development environment
 - C, C++, OpenCL, & DirectCompute platforms
 - DirectX and OpenGL graphics APIs

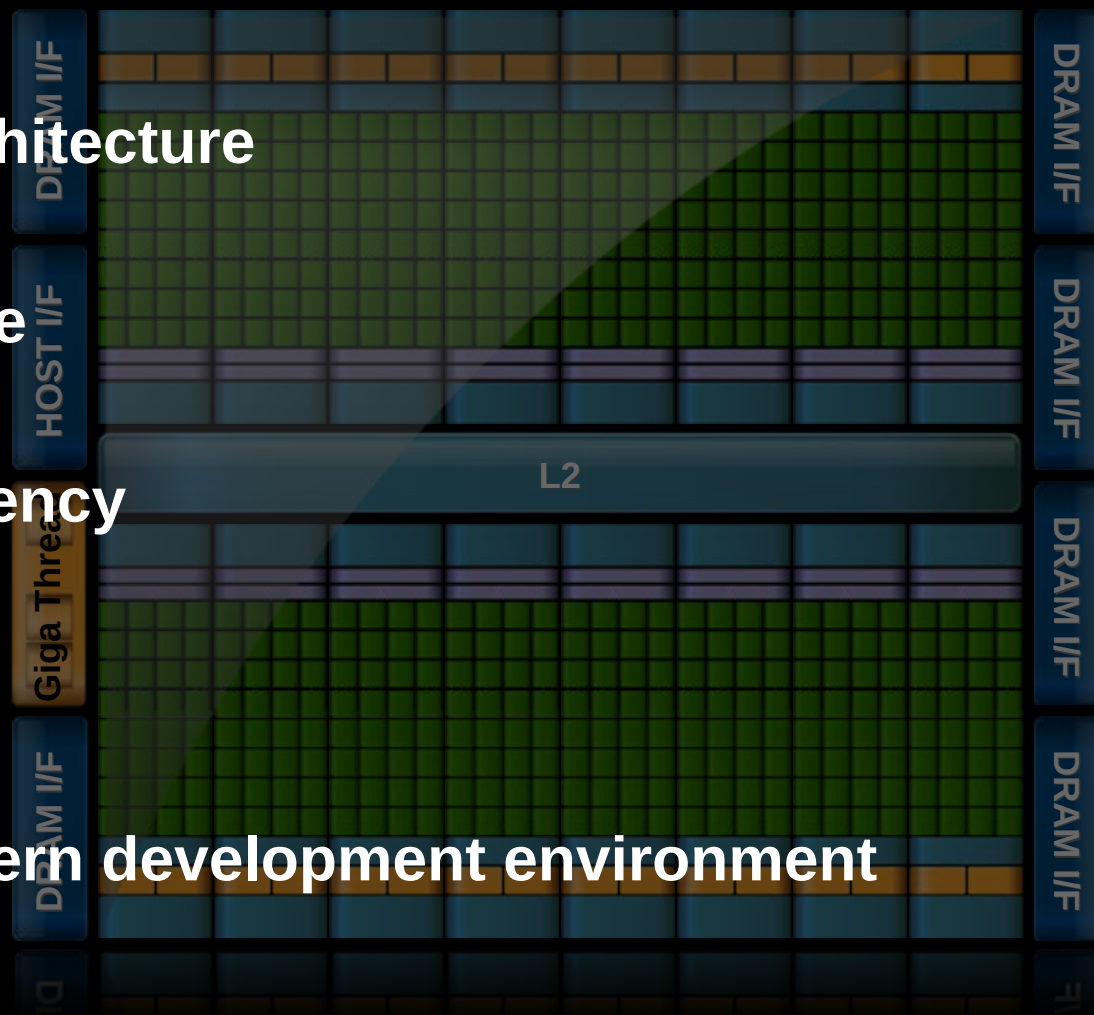




Fermi Summary



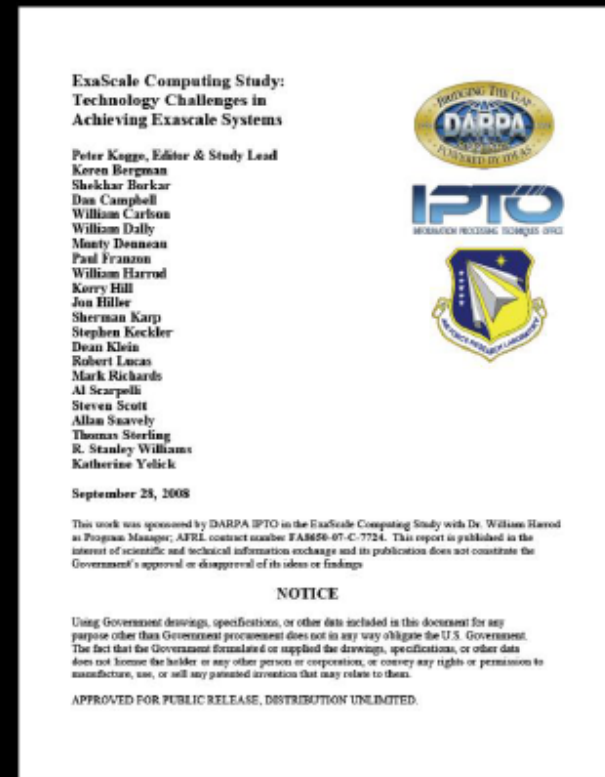
- Third generation CUDA architecture
- Improved peak performance
- Improved throughput efficiency
- Broader applicability
- Full integration within modern development environment



Future: ExaScale Computing



- ExaScale (10^{18}) is the next milestone following PetaScale (10^{15}) computing
- DARPA, HP, France and others are planning for ExaScale computing facilities.
- DARPA report (28th Sept., 2008) described four major challenges (Power consumption, memory, Concurrency/Locality, Resiliency)
- DARPA reports that the number one issue is Power. Extrapolation of Power indicates over 100MW for Exaflop.



Available at
www.darpa.mil/personnel/docs/ExaScale_Study_Initial.pdf

A Possible NVIDIA ExaScale Machine in 2017

by Bill Dally: Chief Scientist & Sr. VP of Research, NVIDIA

A projection based on Moore's Law and does not represent a committed roadmap

- **GPU Node (GPU + CPU + memory + supply) ~300W**
 - 2,400 throughput-cores (7,200FPUs), 16 CPUs – single chip
 - 40TFLOPS (SP), 13TFLOPS (DP)
- **Node Memory**
 - 128GB DRAM, 2TB/s bandwidth
 - 512GB Phase-change Flash for checkpoint and scratch
- **Cabinet ~100kW**
 - 384-Nodes 15.7PFLOPS (SP), 50TB DRAM
 - Dragonfly network – 1TB/sec node bandwidth
- **System ~ 10MW**
 - 128 cabinets – 2 ExaFLOPS (SP), 6.8PB DRAM
 - Dragonfly network with active optical links

Conclusion



- **Performance of Single-threaded processor is saturating**
- **Performance = Parallelism**
- **NVIDIA are committed to offer massive parallel GPUs for now and beyond**
- **CUDA GPU computing abstracts parallel programming**
- **'Fermi' is the key milestone in GPU Computing**

Thank you for your attention

