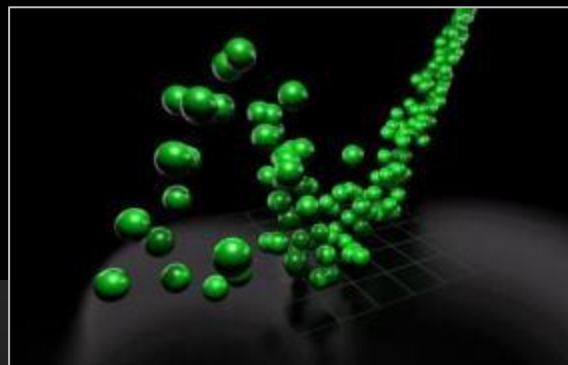
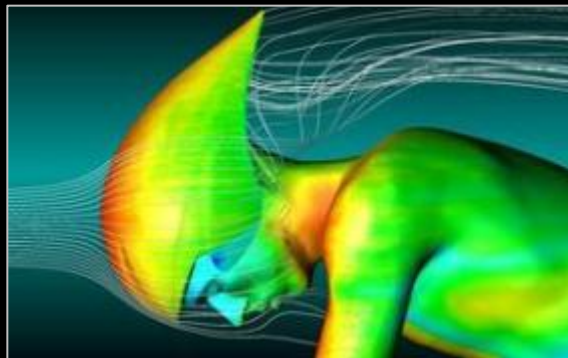


NVIDIA



Atelier Teratec 2011 'Etat de l'art GPU Computing'



1995

5,000 triangles/second
800,000 transistors GPU





2011

350 Million triangles/second
3 Billion transistors GPU



NVIDIA
SuperPhones to SuperComputers

GPU Computing Today

300 Million

CUDA Capable GPUs

1 Million

CUDA Toolkit Downloads

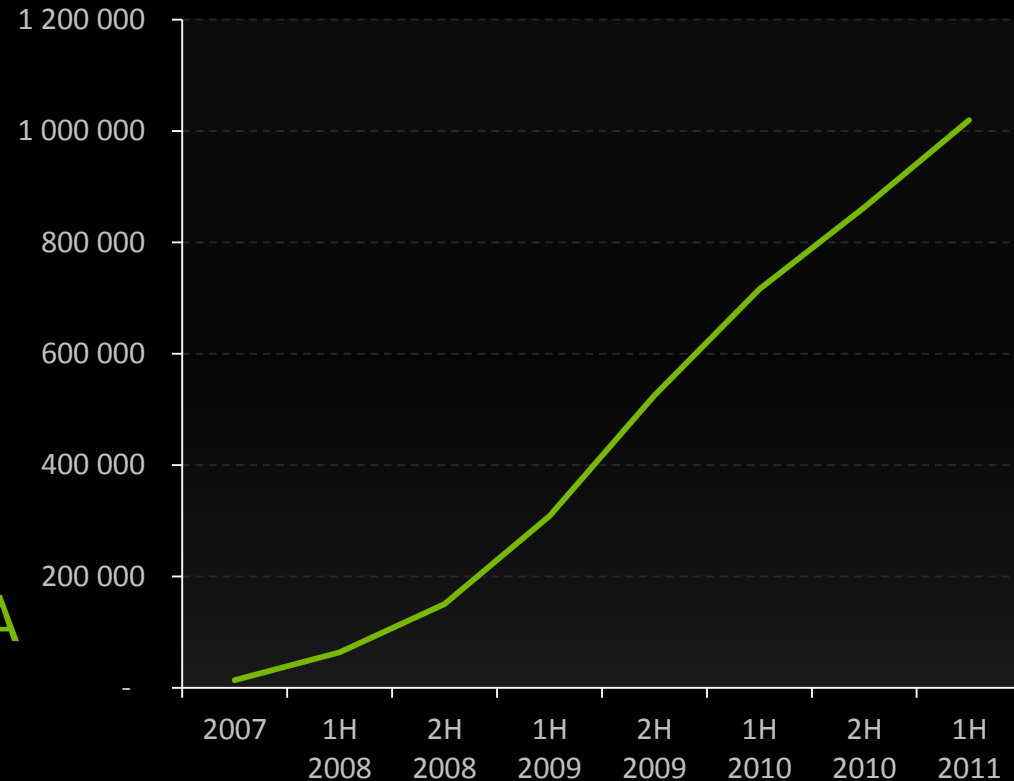
150 Thousand

Active CUDA Developers

450

Universities Teaching CUDA

1 Million CUDA
Developer Toolkit Downloads



Getting Started with GPUs

TRY

Try CUDA Toolkit on your Notebook or Desktop with CUDA Enabled GPU



DEVELOP

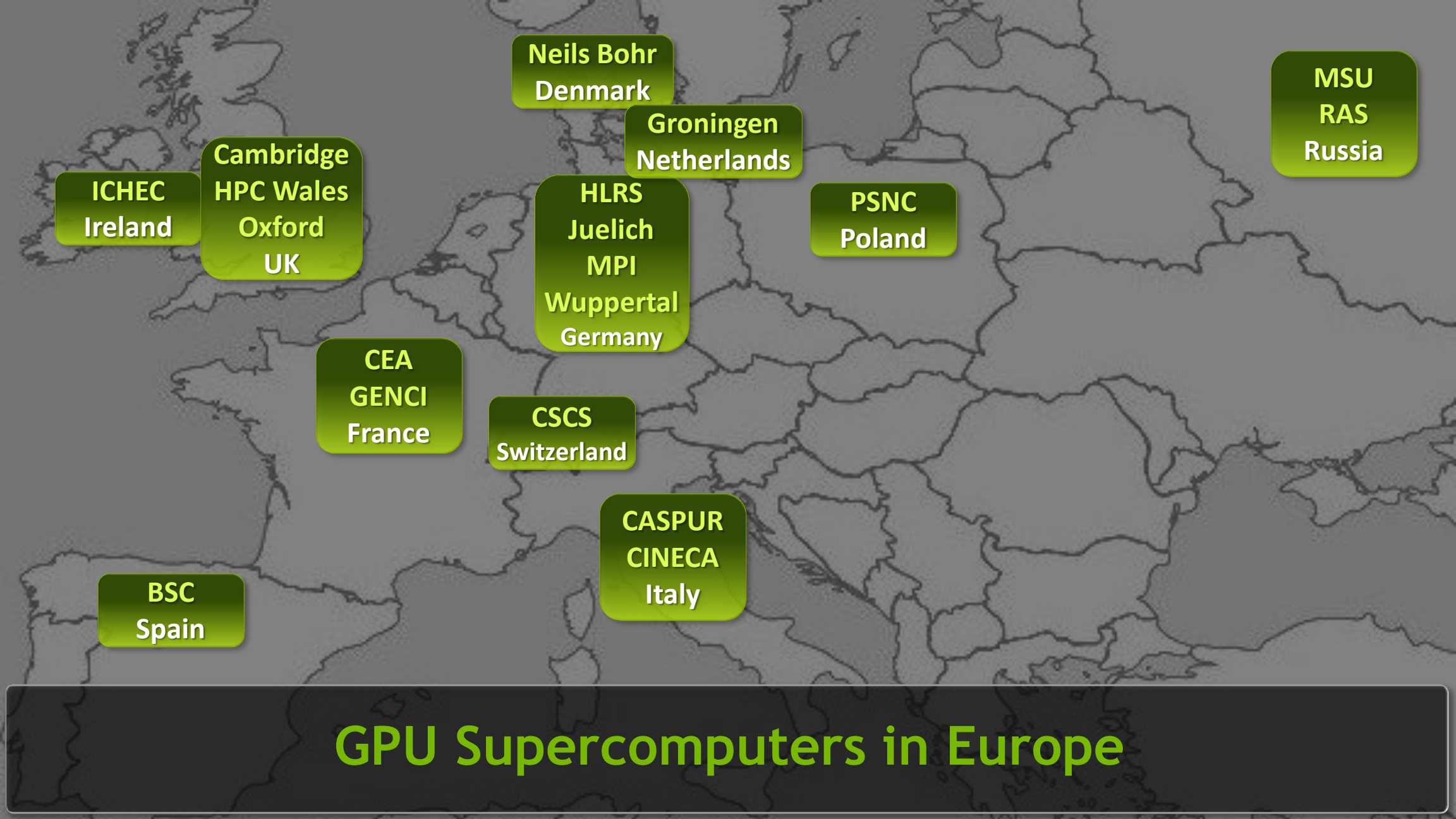
Optimize HPC Apps on Compute Workstation with Tesla GPUs



DEPLOY

Run apps in production with GPU Compute Cluster



A map of Europe with various countries outlined in black. Overlaid on the map are ten green rounded rectangular callouts, each containing the name of a GPU supercomputer project and its location. The callouts are positioned as follows: Ireland (top left), UK (top left, below Ireland), Denmark (top center), Netherlands (top center, below Denmark), Russia (top right), Poland (center right), Germany (center), France (center left), Switzerland (center, below Germany), Italy (center, below Switzerland), and Spain (bottom left).

Neils Bohr
Denmark

Groningen
Netherlands

MSU
RAS
Russia

ICHEC
Ireland

Cambridge
HPC Wales
Oxford
UK

PSNC
Poland

HLRS
Juelich
MPI
Wuppertal
Germany

CEA
GENCI
France

CSCS
Switzerland

BSC
Spain

CASPUR
CINECA
Italy

GPU Supercomputers in Europe

GPU Computing Developer Ecosystem

Numerical Packages

MATLAB/Jacket
Mathematica
NI LabView
pyCUDA

Debuggers & Profilers

cuda-gdb
NV Visual Profiler
Nsight for MS VS
Allinea
TotalView

GPU Compilers

C
C++
Fortran
OpenCL
DirectCompute
Java
Python

Parallelizing Compilers

PGI Accelerator
CAPS HMPP
Par4All
mCUDA

Cluster Management

nv-smi
Bright Cluster
Rocks Roll
Scyid
PBS Pro
Cluster Manager

Libraries

BLAS
FFT
LAPACK
NPP
Video
Imaging

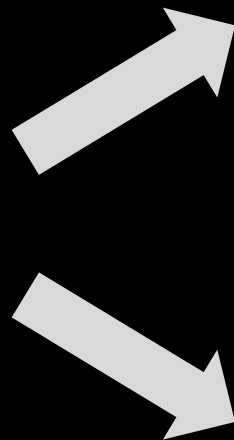
Consultants & Training



Solution Providers



CUDA GPUs: Support All Languages



CUDA C

CUDA C++

CUDA
Fortran

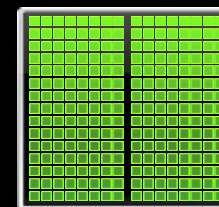
OpenCL

Direct-
Compute

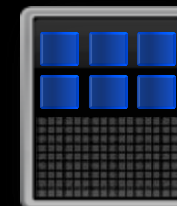
C++ AMP



**NVIDIA
CUDA GPUs**



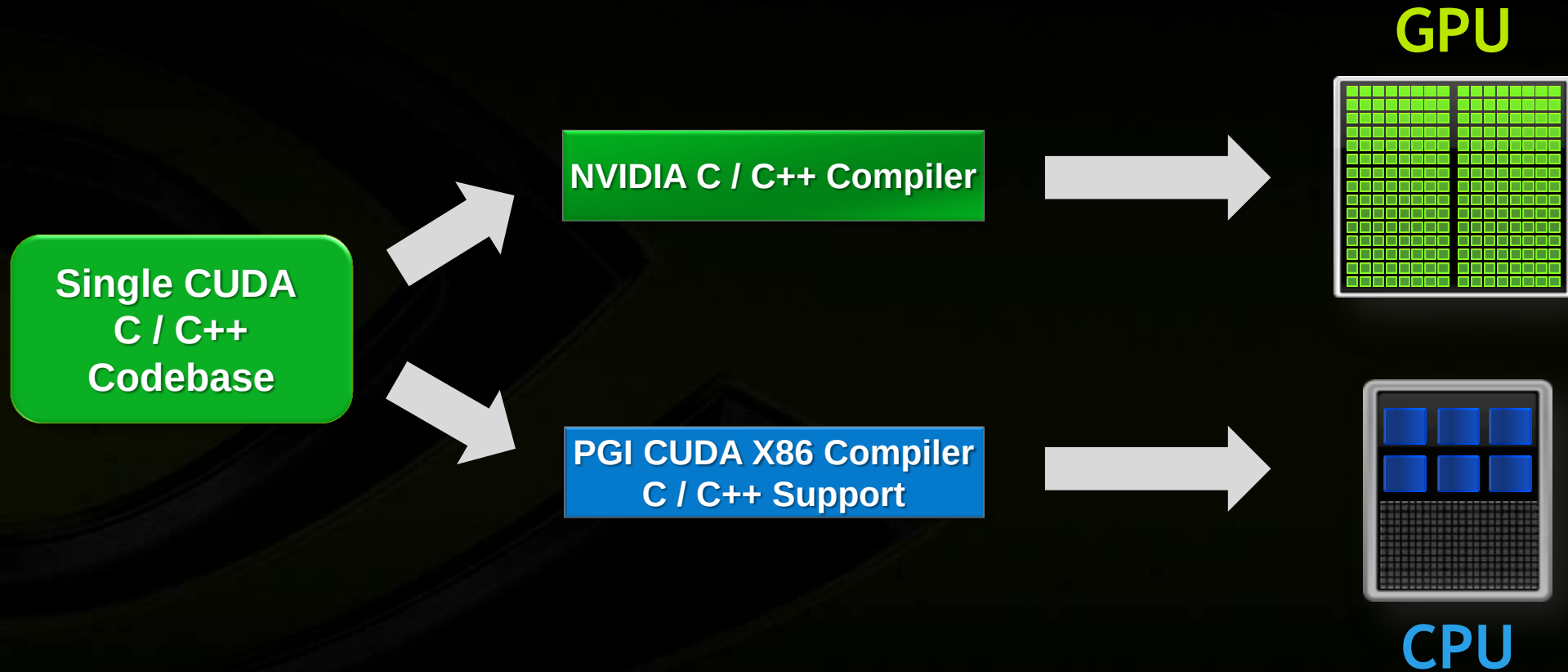
CPU



**Other
Devices**

PGI CUDA x86

CUDA Now Available for CPUs and GPUs



```

1 integer, parameter :: n = 1000
2
3 real, dimension(n) :: input
4 real, dimension(n) :: result
5 integer :: i
6 do i = 1,n
7 input(i) = i*4.0
8 enddo
9 !$omp acc_region
10 !$omp acc loop
11 do i = 1,n
12 result(i) = input(i) * 4.0
13 enddo

```

On this line:
1 Process: rank 0

CRAY
THE SUPERCOMPUTER COMPANY

U threads
>>> ... <<<(0,0),(0,0,0)>>> (32 threads)

```

1 # 29 "basic.c"
2 void hmp_codelet_myFunc(int n, int A[n], const int B[n])
3 {
4     #pragma hmpccg gridify(i)
5     # 7 "<preprocessor>"
6     # 32 "basic.c"
7     for (int i = 0 ; i < n ; ++i)
8     {
9         A[i] += B[i] + 1;
10    }
11 }
12
13
14
15
16
17

```

On this line:
1 Process: rank 0

Kernel 1: 32 GPU threads
<<<(0,0),(0,0,0)>>> ... <<<(0,0),(31,0,0)>>> (32 threads)

CAPS
Innovative software for manycore paradigms

- Supporting the environments that you will use for hybrid development

```

1 #pragma acc region
2 {
3     for( i = 0; i < n; ++i ) r[i] = a[i]*2.0f;
4 }
5
6 /* compute on the host to compare */
7 for( i = 0; i < n; ++i ) e[i] = a[i]
8 < the results */
9 = 0; i < n; ++i )
10 prt( r[i] == e[i] ); . . .

```

On this line:
1 Process: rank 0
1 Thread (Process 0): #1

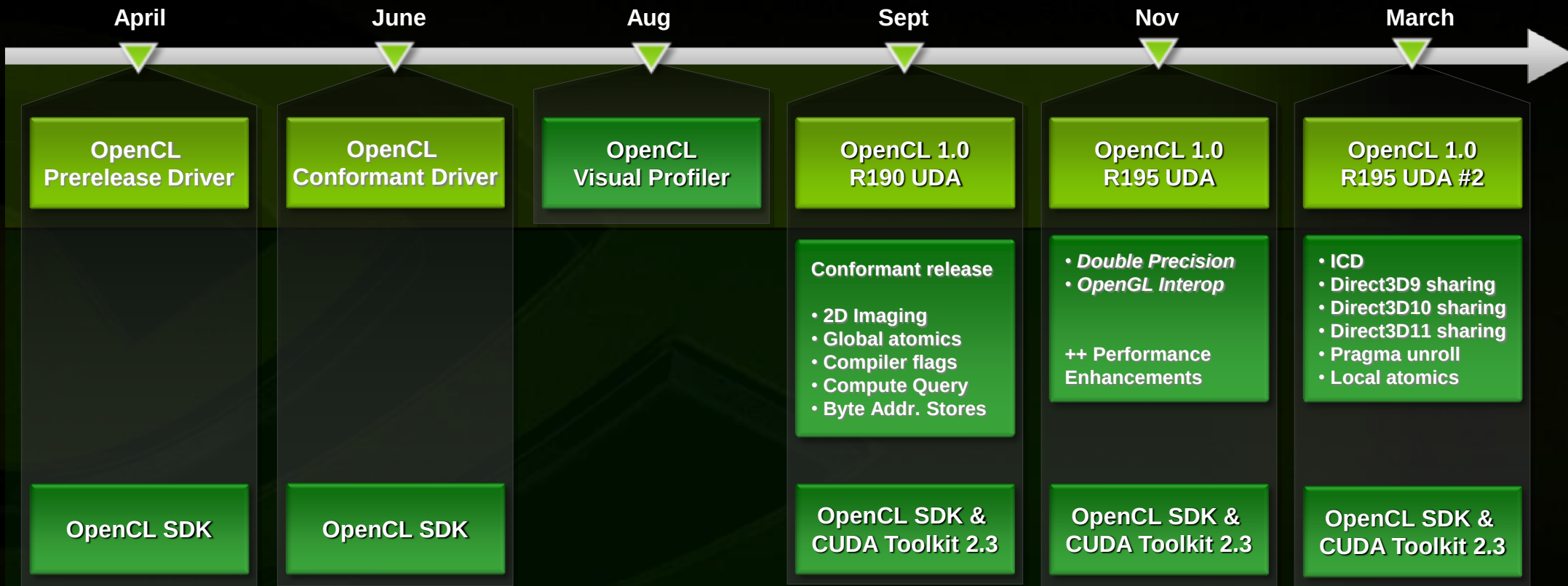
PGI

NVIDIA OpenCL Execution



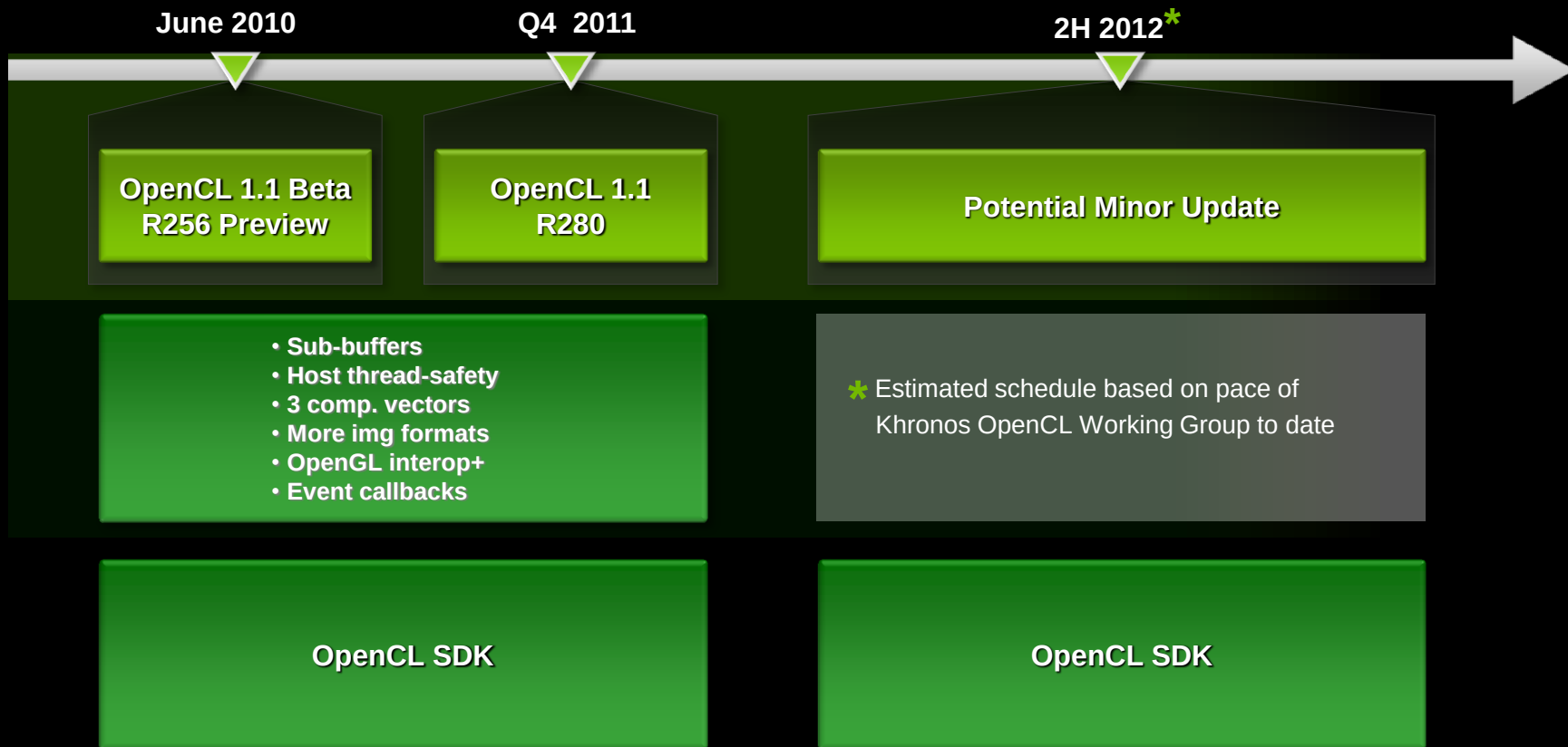
2009

2010





NVIDIA OpenCL Roadmap



CUDA 4.0: Big Leap In Usability



Ease
of Use

CUDA 1.0
Program GPUs using
C, Libraries

CUDA 2.0
Debugging, Profiling,
Double Precision

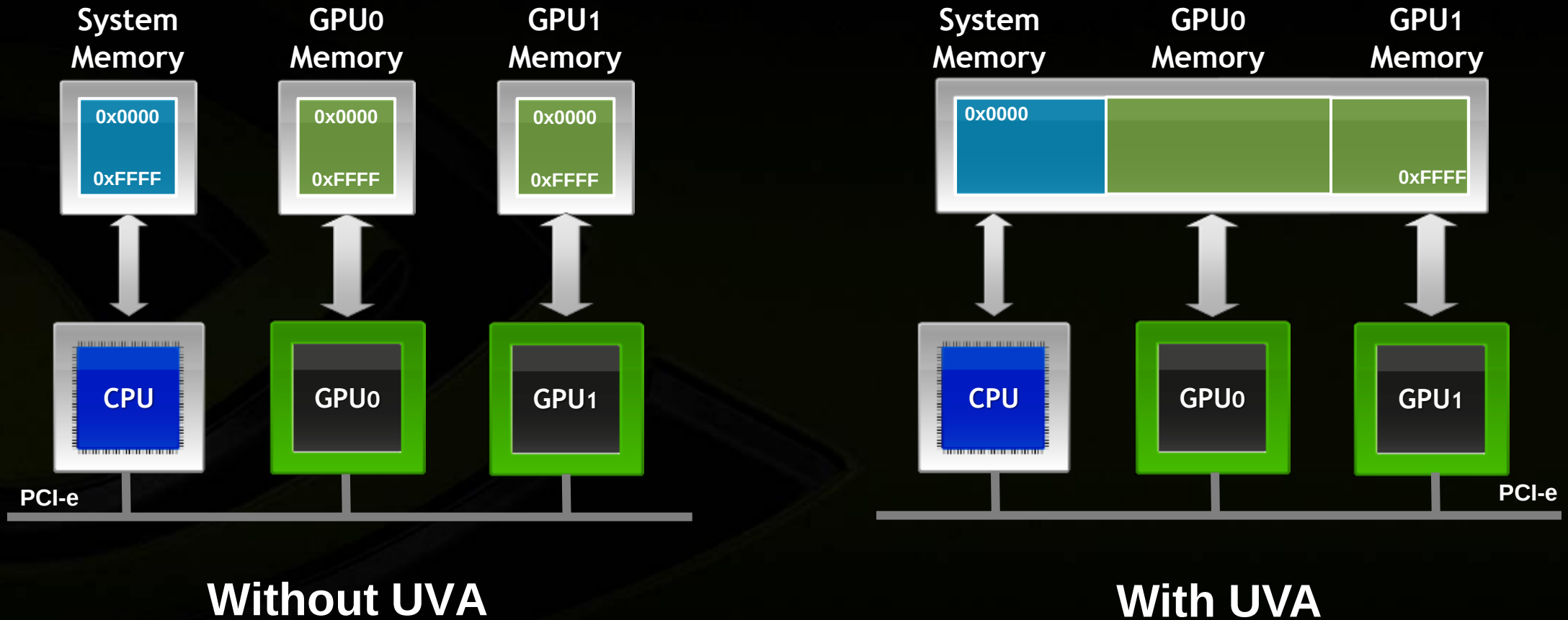
CUDA 3.0
Fermi, New Libraries,
Big Perf. Boost

CUDA 4.0
Parallel Programming
Made Easy

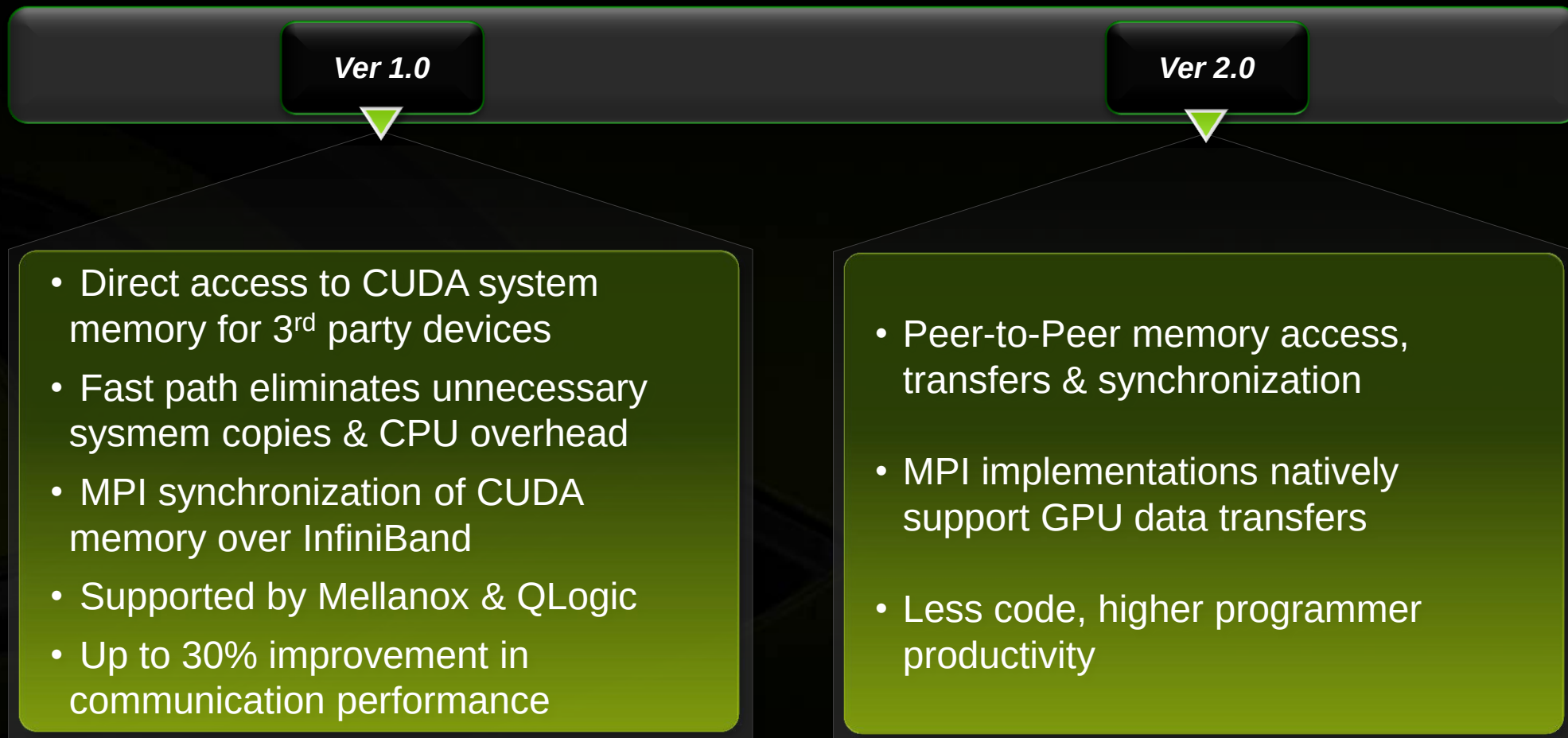
Performance

Unified Virtual Addressing

Easier to Program with Single Address Space

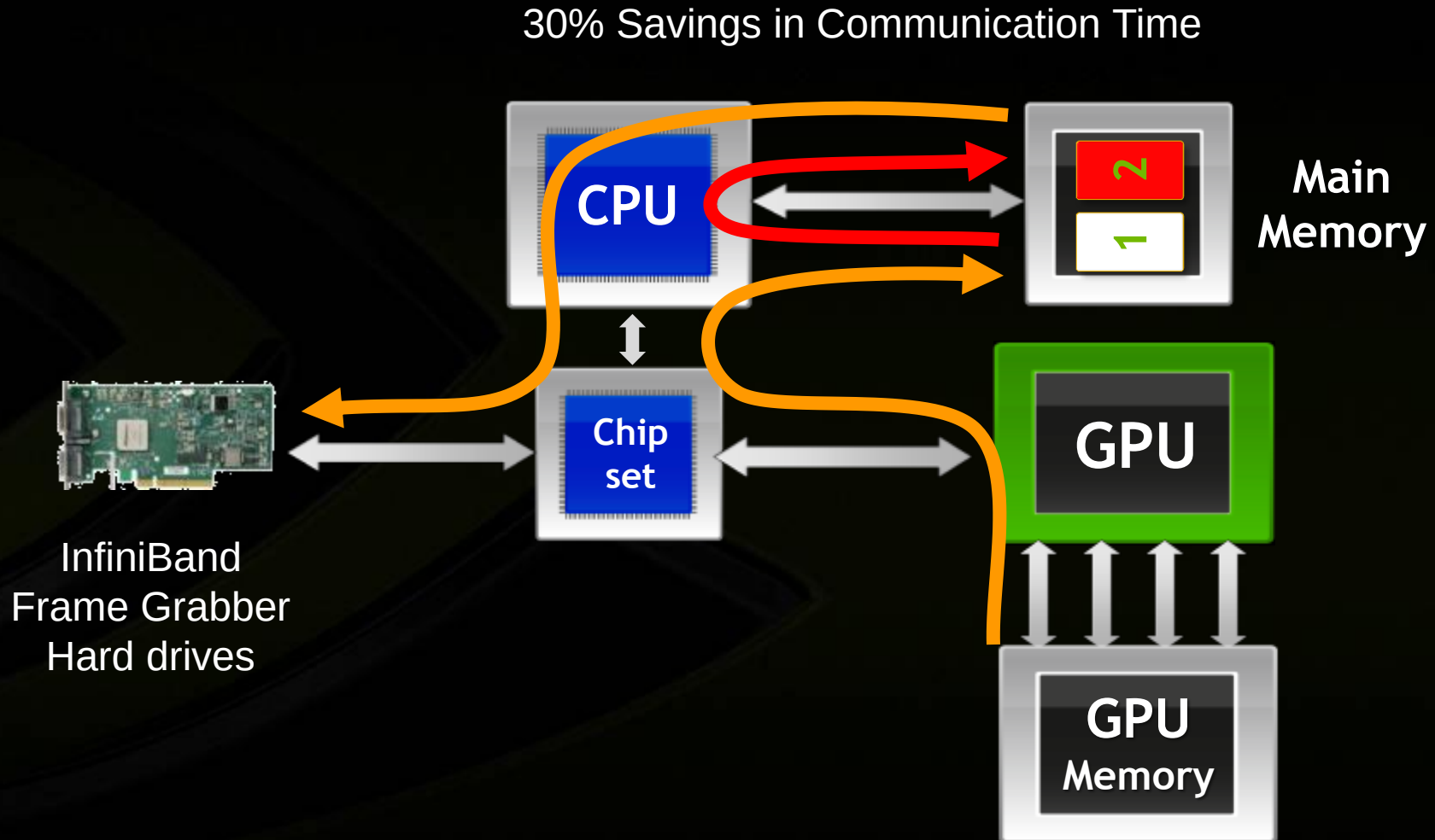


NVIDIA GPUDirect™

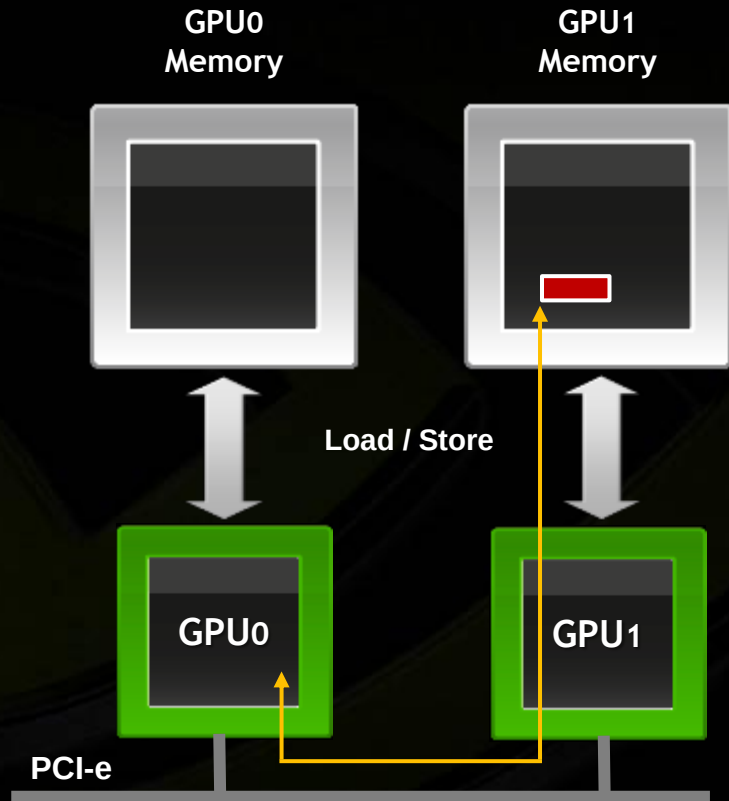


Details @ <http://developer.nvidia.com/object/gpudirect.html>

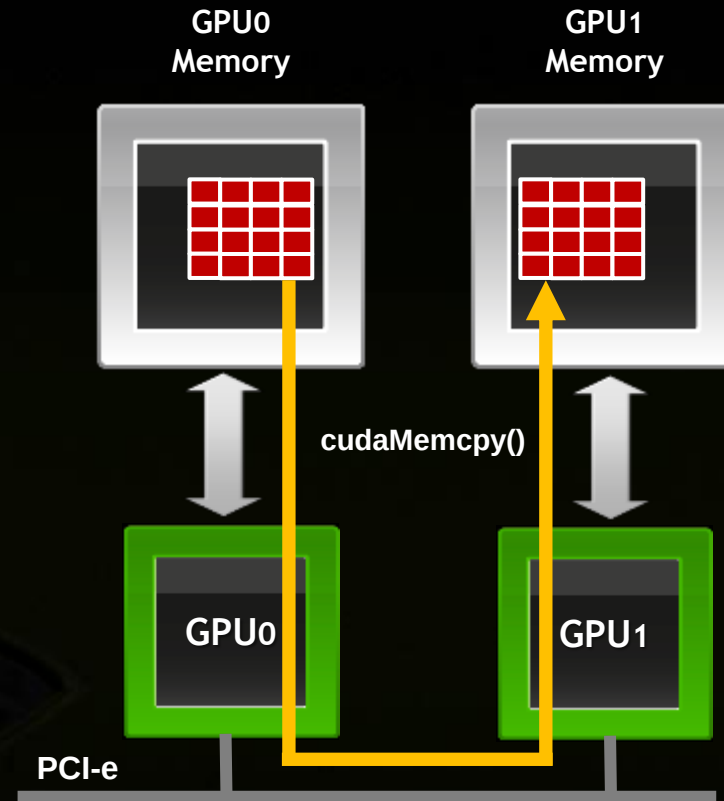
NVIDIA GPUDirect 1.0: Faster Communication with GPUs



GPUDirect v2.0: Peer-to-Peer Communication

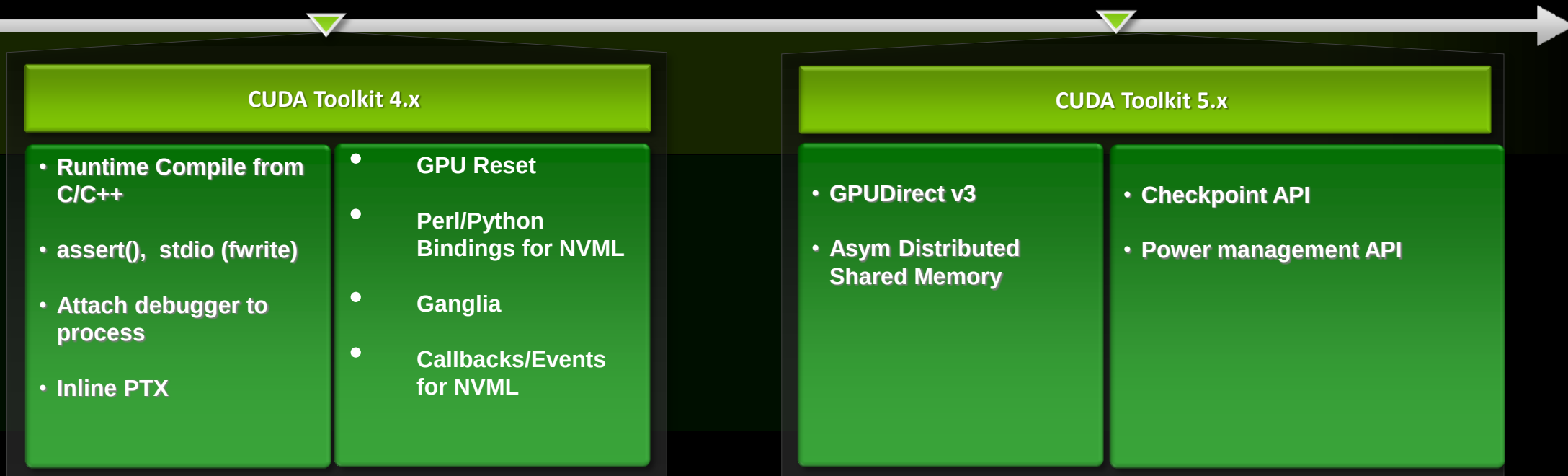


Direct Access



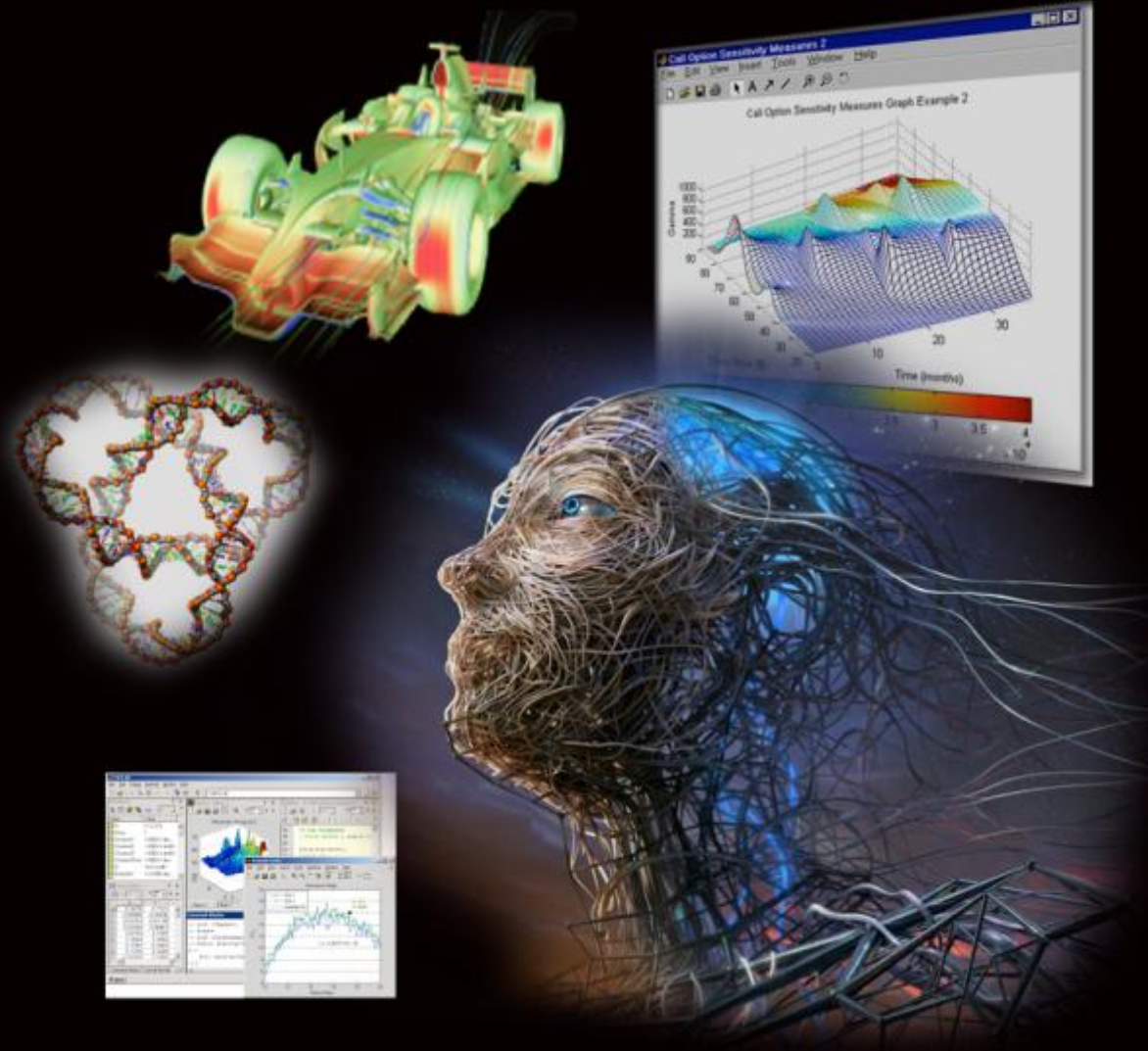
Direct Transfers

CUDA Roadmap



Strategic Focus on Applications

- Senior-level relationship and market managers
- Dedicated technical resources
- More than 150 people devoted to libraries, tools, application porting and market development
- Worldwide focus



Wide Adoption of Tesla GPUs

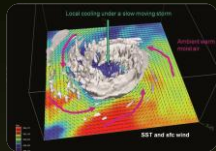
Oil and gas



Reverse Time
Migration
Kirchoff Time
Migration
Reservoir Sim



Edu/Research



Astrophysics
Molecular
Dynamics
Weather / Climate
Modeling



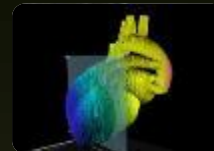
Government



Signal Processing
Satellite Imaging
Video Analytics
Synthetic Aperture
Radar



Life Sciences



Bio-chemistry
Bio-informatics
Material Science
Sequence Analysis
Genomics



Finance



Risk Analytics
Monte Carlo
Options Pricing
Insurance
modeling



Manufacturing



Structural
Mechanics
Computational
Fluid Dynamics
Machine Vision
Electromagnetics



Strategic Partners



CAD/ CAM/ CAID	CAE/ EDA	Computational chemistry	Computational Finance	Defence & Intelligence	Digital Content creation	Physical Sciences	Seismic processing and visualization
Autodesk	Ansys	Amber	MATLAB	Ikena	Adobe	Quda (L-QCD)	Schlumberger
Dassault Systemes: CATIA Solidworks	Dassault Systemes: Simulia	NAMD	Mathematica	Intergraph	Autodesk M&E	WRF	Landmark
PTC	Nastran	Gromacs	NAG	ESRI	Avid	ACUSA	Paradigm
Siemens	LSTC	Lammps	Murex	Manifold	MainConcept	HOMME	
	Synopsys	GAMESS			Sony	HYCOM	

CUDA GPU Roadmap



NVIDIA Tesla M2070-Q



With full Quadro OpenGL

- Quadro Enhanced performance profile
- Quadro Optimized profiles for professional applications
- Certified applications
- Remote OGL Visualization Support
- 3rd Party Remote Support: MS RemoteFX, Citrix HDX
- AXE – Application Acceleration Engines
- Multi GPU: SLI Multi- OS, SLI Mosaic



Same dimensions, same thermal as Tesla M2070
Dual slot width, no video out connector

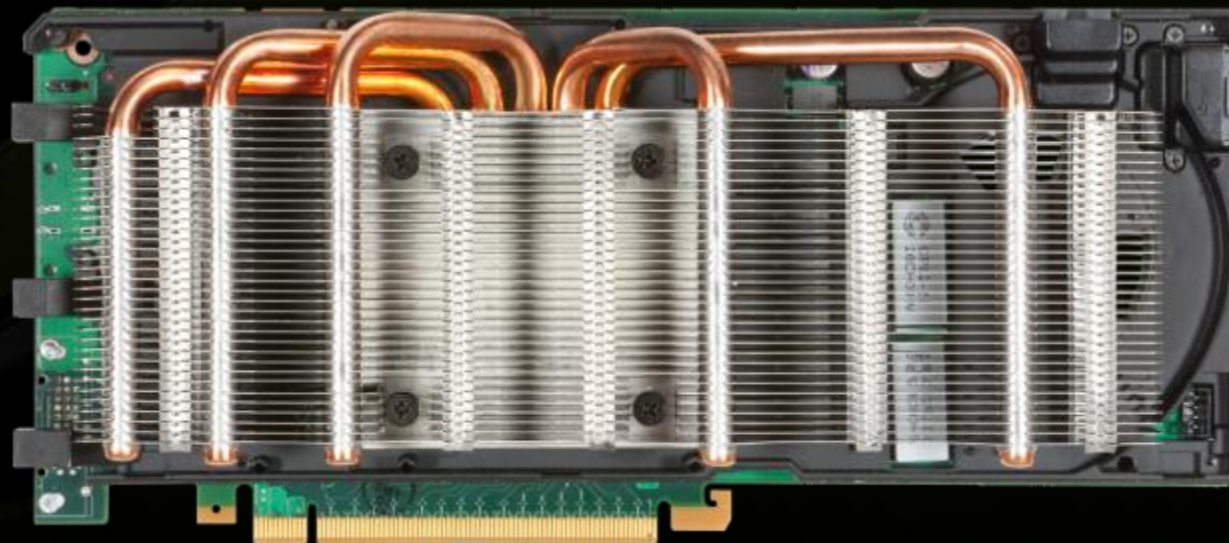
40% Smaller Tesla Module



Tesla X



Tesla M-series



Fastest Supercomputer in Russia



Moscow State University

1.3 Petaflops Peak

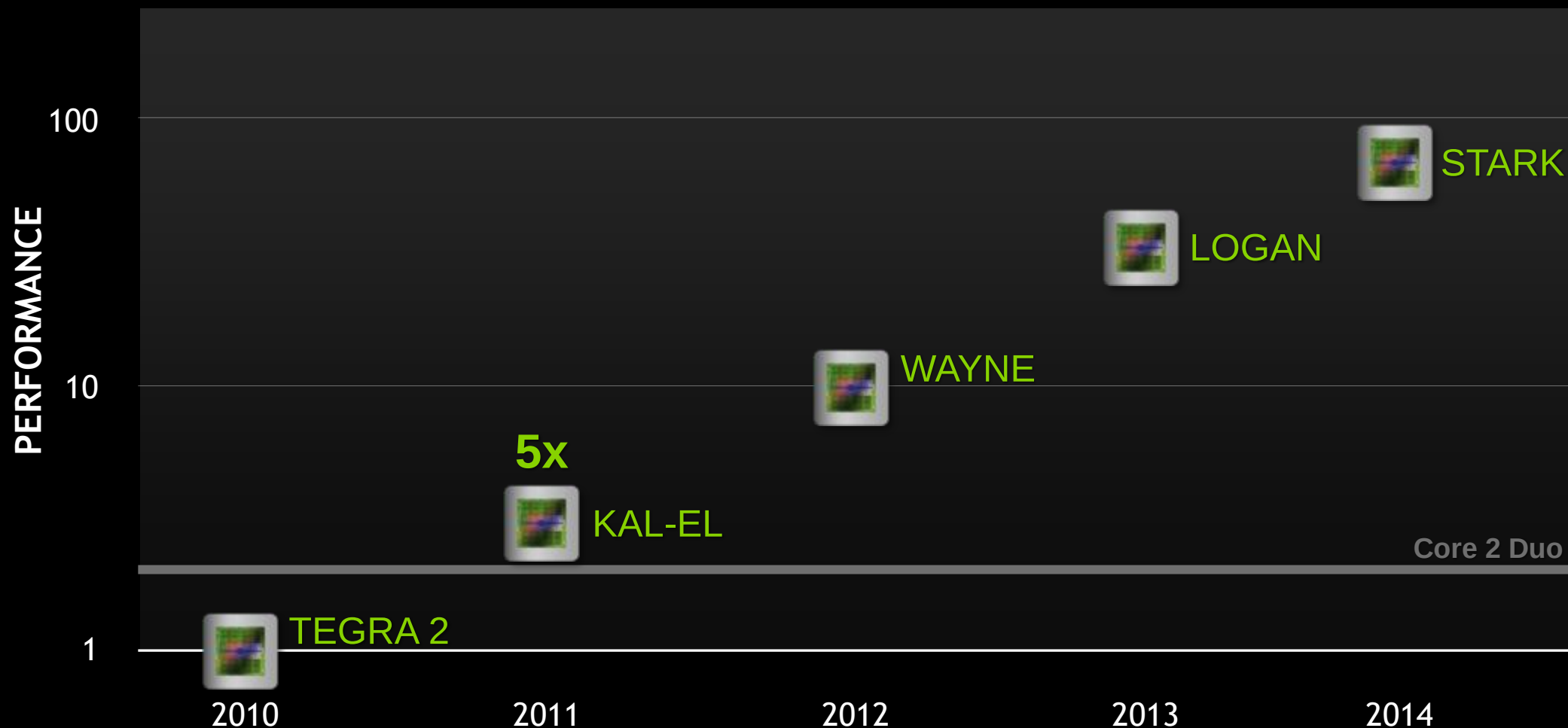
1554 Tesla X2070 GPUs

+1554 Intel CPUs

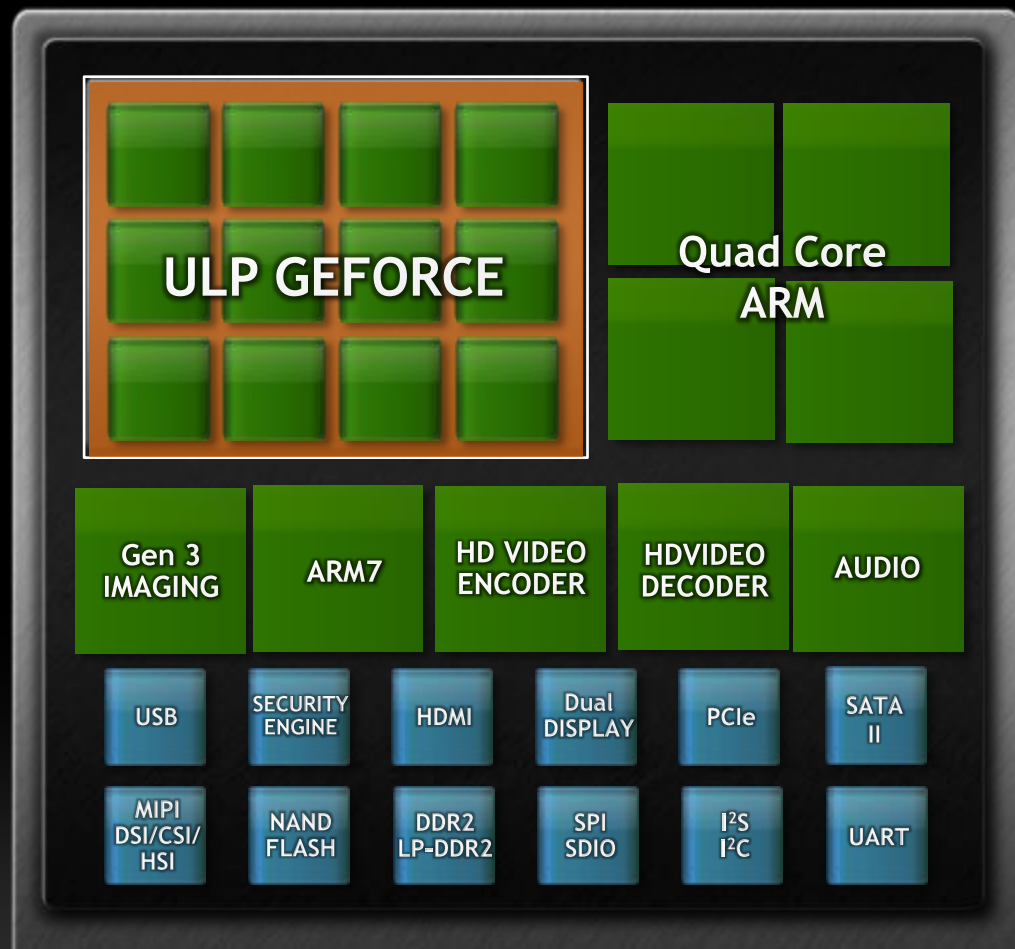
Research Topics

- **Ocean Modeling**
- **Climate Change**
- **Genomic Medicine**
- **Galaxy Formation**

Tegra Roadmap

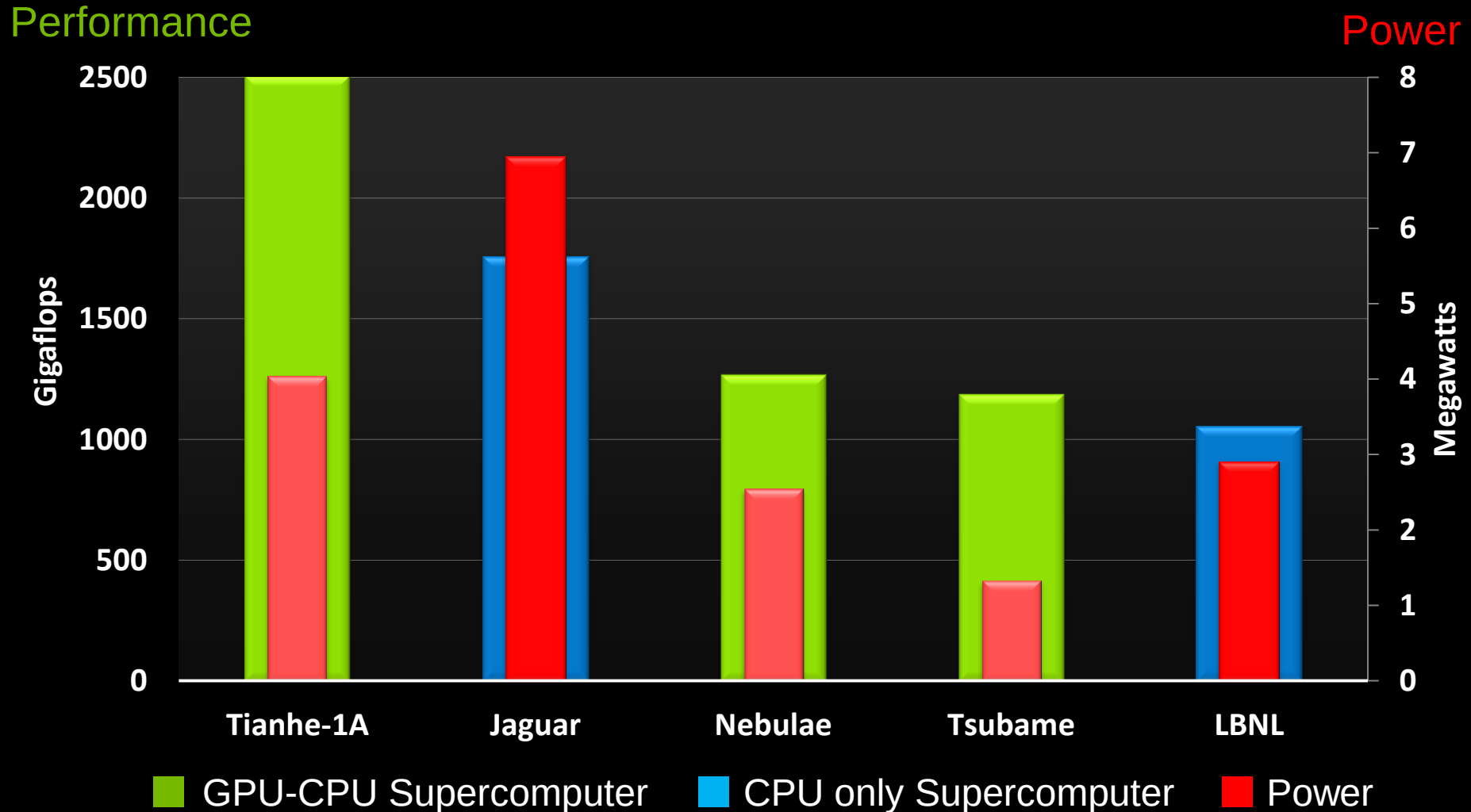


Kal-El — Independent, Dedicated Processing



CPU	3X Performance <i>Quad Core up to 1.5GHz, NEON</i>
POWER	20x Lower Power <i>Due to ULP mode</i>
VIDEO	4X Complexity <i>1080i/p High Profile</i>
GRAPHICS	3X Performance <i>12 Core, Dual Pixel Pipe</i>
MEMORY	3X bandwidth <i>DDR3L up to 1600 data rate</i>
IMAGING	Better noise reduction & color rendition <i>Two simultaneous streams</i>
AUDIO	HD Audio <i>7.1 channel surround</i>
STORAGE	2 - 6X faster <i>e.MMC 4.4 and SATA-II</i>

GPU Supercomputers: More Power Efficient



To ExaScale and Beyond

DARPA Study Identifies Four Challenges for ExaScale Computing



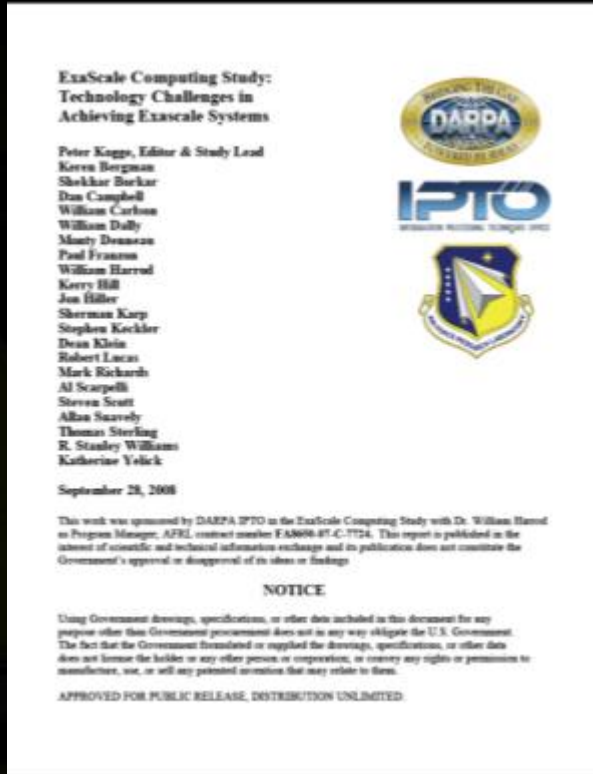
Report published September 28, 2008:

- **Four Major Challenges**

- Energy and Power challenge
- Memory and Storage challenge
- Concurrency and Locality challenge
- Resiliency challenge

- **Number one issue is *power***

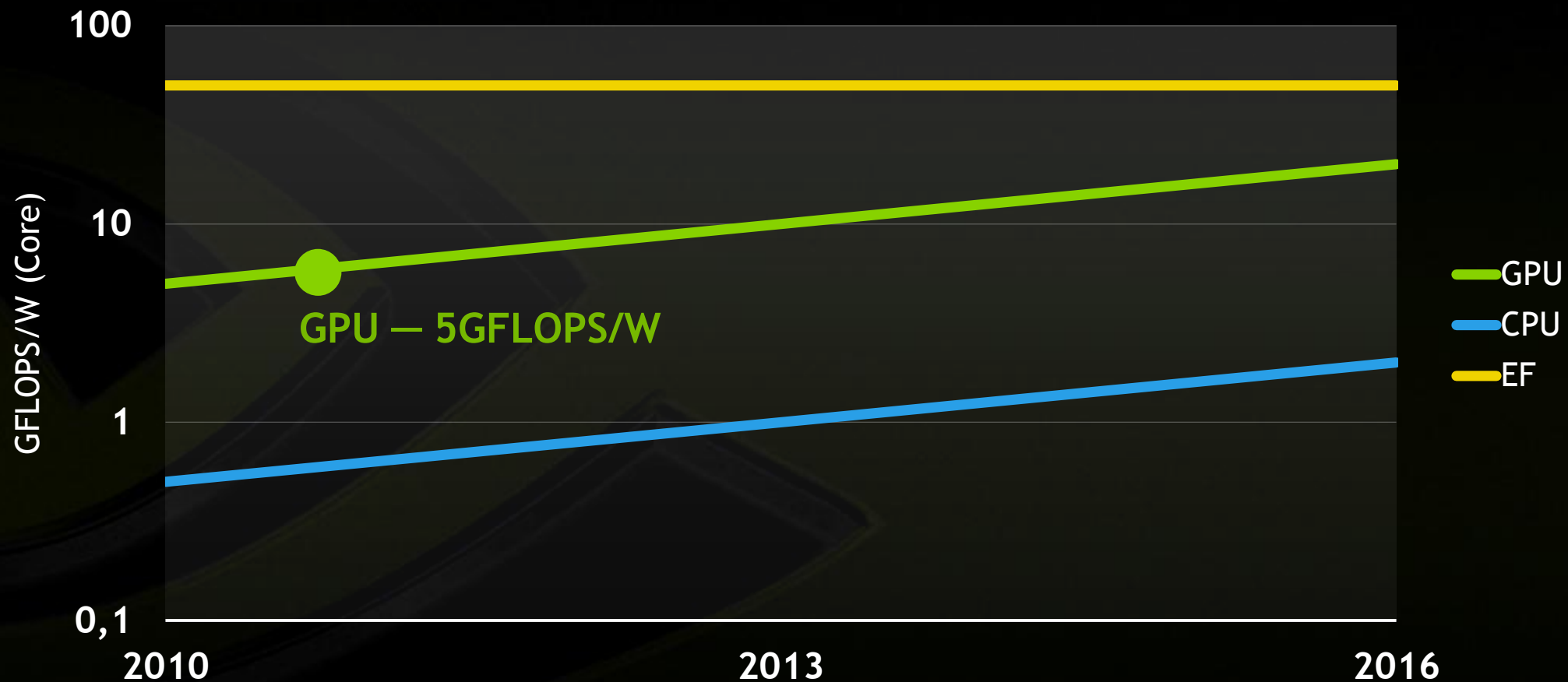
- Extrapolations of current architectures and technology indicate over 100MW for an Exaflop!
- Power also constrains what we can put on a chip



Available at

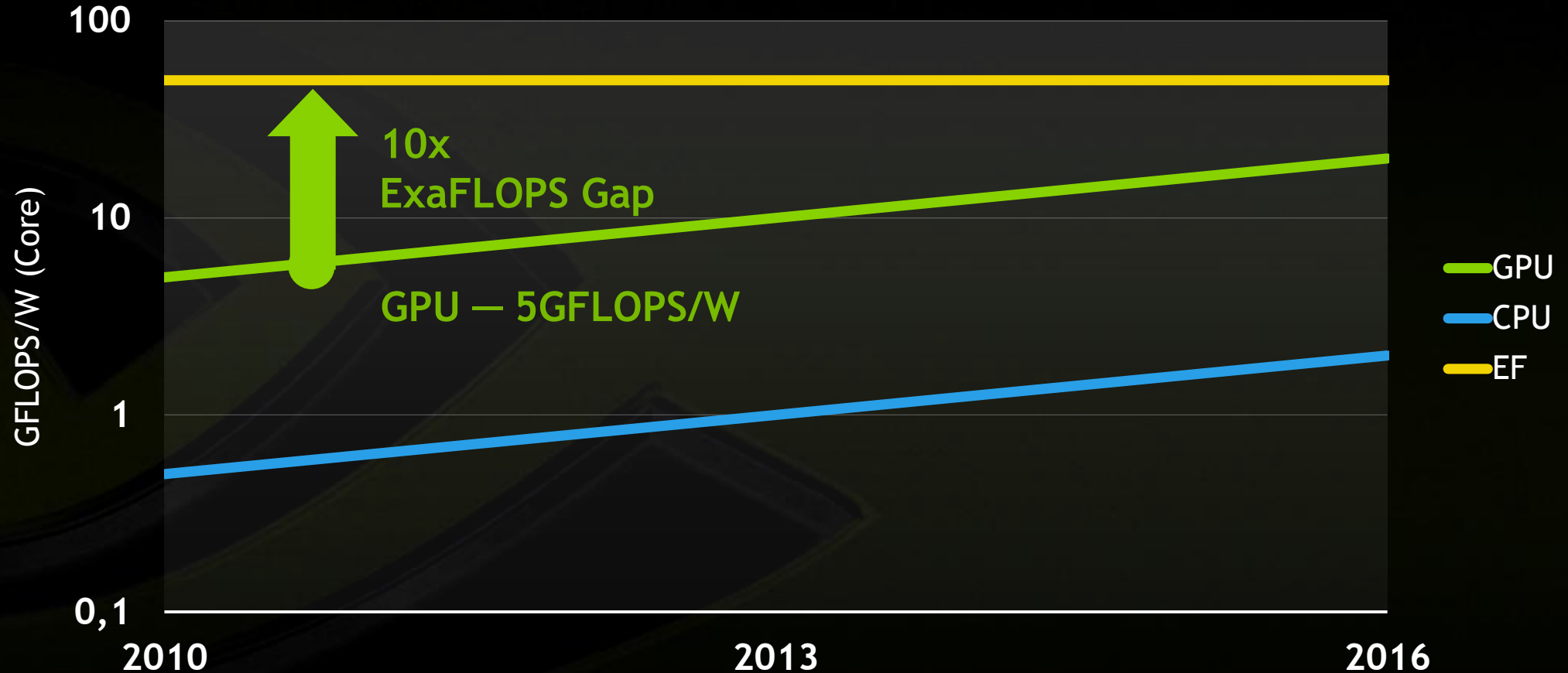
www.darpa.mil/ipto/personnel/docs/ExaScale_Study_Initial.pdf

ExaFLOPS at 20MW = 50GFLOPS/W



50GFLOPS/W

10x Energy Gap for Today's GPU



Echelon

NVIDIA's Extreme-Scale Computing Project

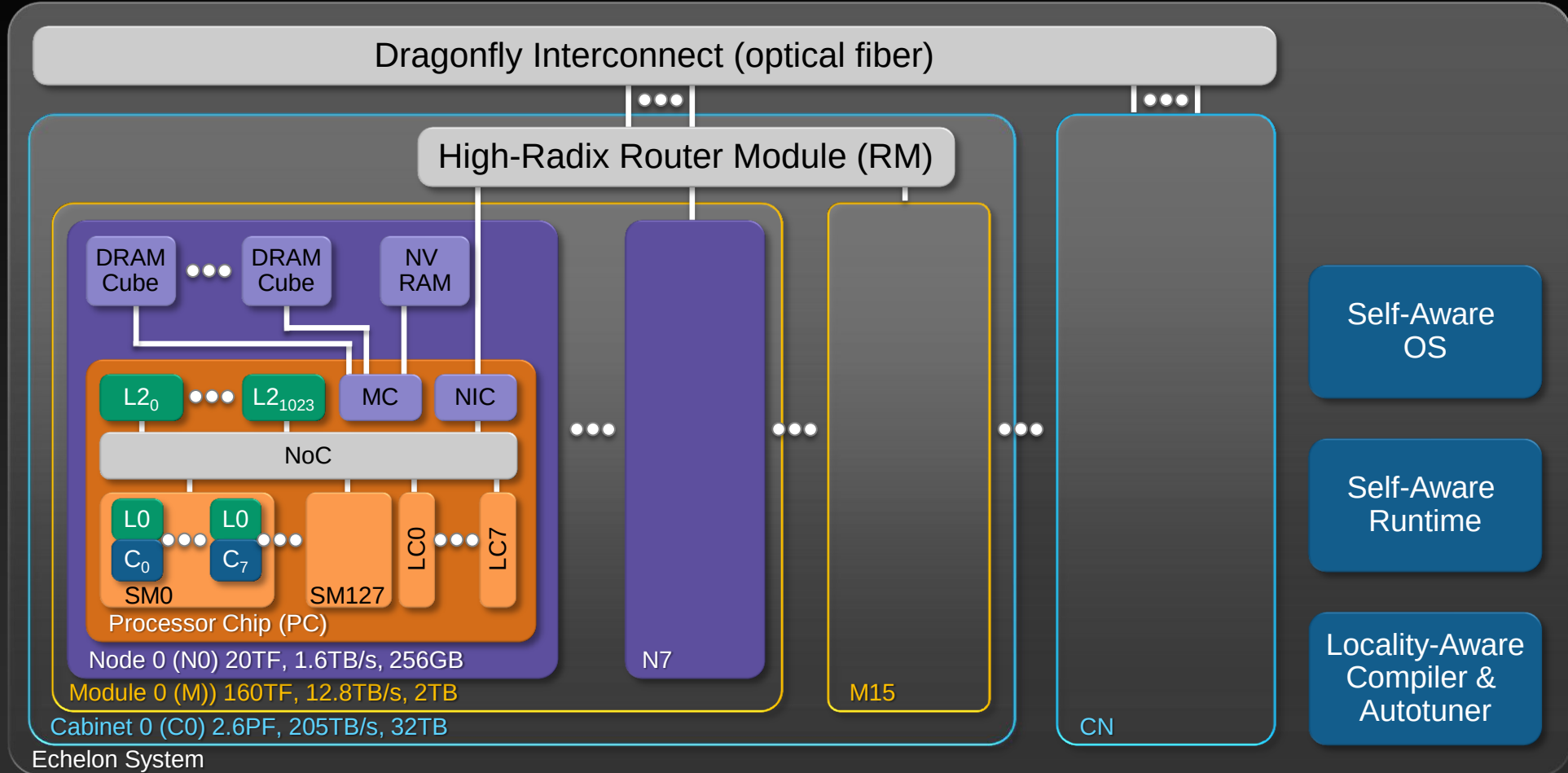
Echelon Team



LOCKHEED MARTIN



System Sketch



Jean-Christophe Baratault

ibaratault@nvidia.com

Cell +33 6 8036 8483

