

TSUBAME2.0: A Tiny and Greenest Petaflops Supercomputer

Satoshi Matsuoka

Global Scientific Information and Computing
Center (GSIC)

Tokyo Institute of Technology (Tokyo Tech.)
Booth #1127

Booth Presentations SC10

Nov 2010

The TSUBAME 1.0 "Supercomputing Grid Cluster"

Spring 2006

Unified IB network

Voltaire ISR9288 Infiniband 10Gbps

x2 (DDR next ver.)

~1310+50 Ports

~13.5Terabits/s (3Tbits bisection)



*"Fastest
Supercomputer in
Asia" 7th on the 27th
Top500@38.18TF*

Sun Galaxy 4 (Opteron Dual
core 8-socket)

10480core/655Nodes

21.4Terabytes

50.4TeraFlops

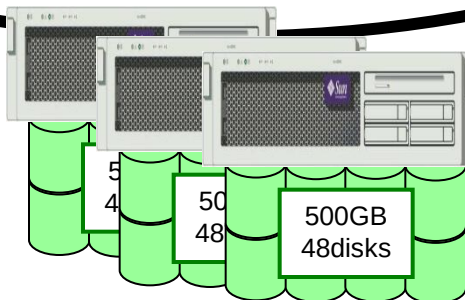
OS Linux (SuSE 9, 10)

NAREGI Grid MW

10Gbps External
Network



NEC SX-8i
(for porting)



Storage

1.0 Petabyte (Sun "Thumper")

0.1Petabyte (NEC iStore)

Lustre FS, NFS, WebDAV (over IP)

50GB/s aggregate I/O BW

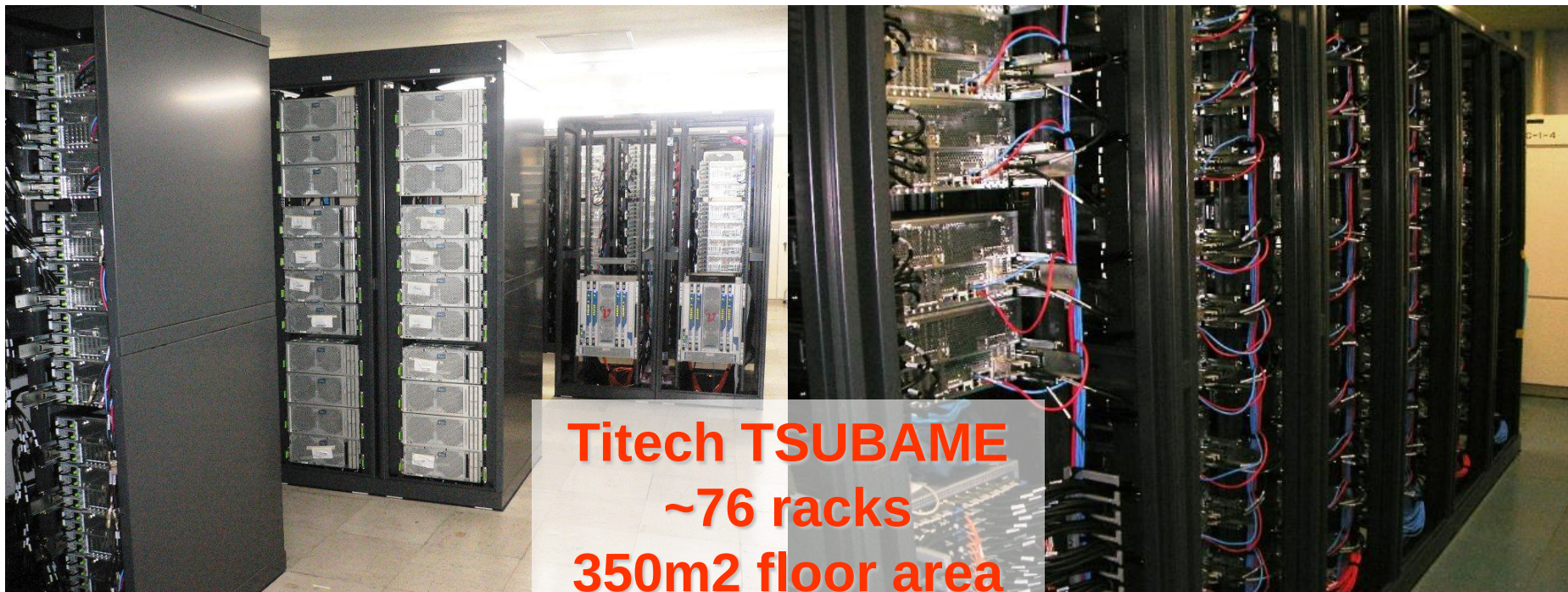


ClearSpeed CSX600

SIMD accelerator

360 boards,

35TeraFlops(Current))

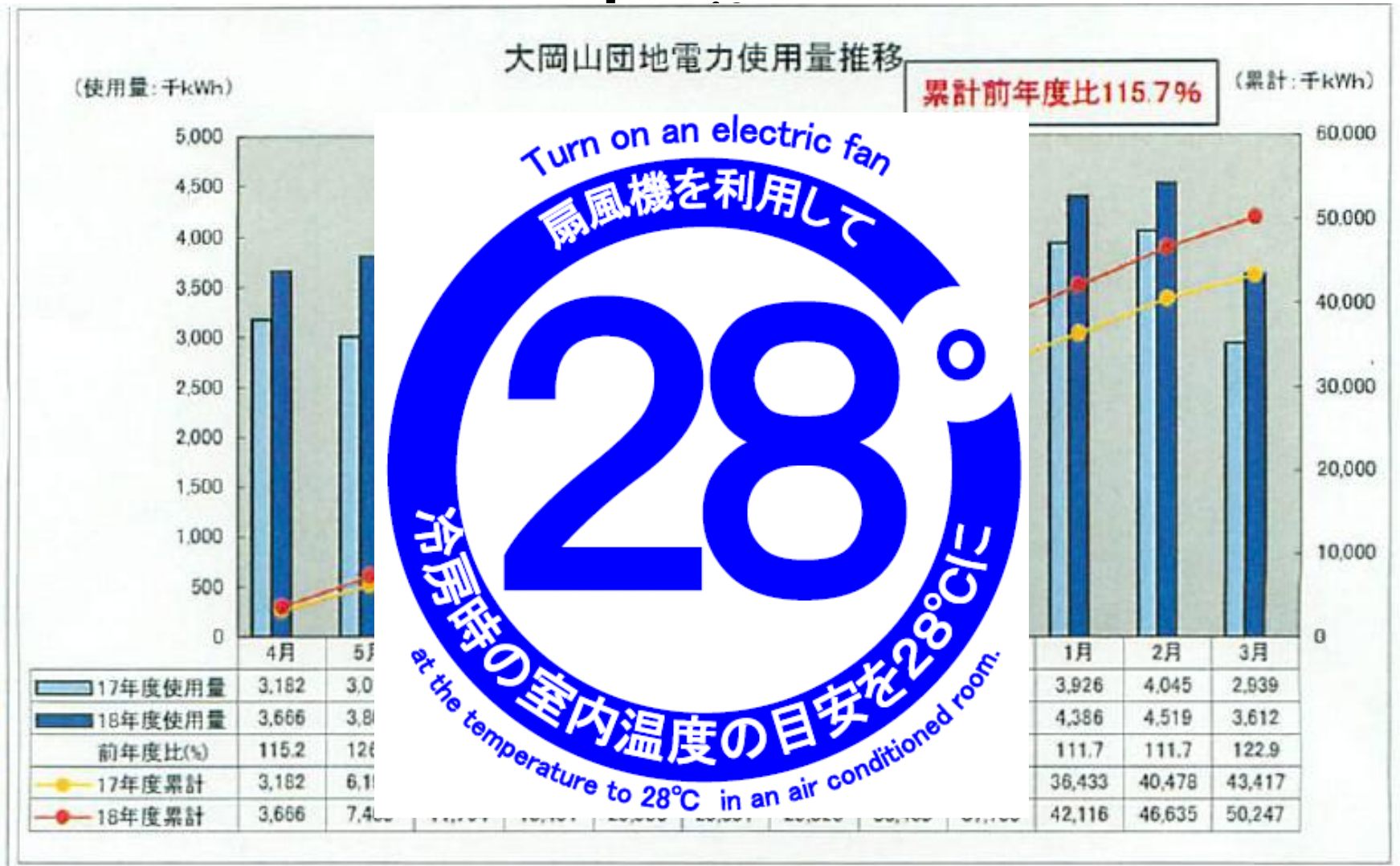


Titech TSUBAME
~76 racks
350m2 floor area
1.2 MW (peak),
PUE=1.44



You know you have a problem

1



Biggest Problem is Power...

Machine	CPU Cores	Watts	Peak GFLOPS	Peak MFLOPS/ Watt	Watts/ CPU Core	Ratio c.f. TSUBAME
TSUBAME(Opteron)	10480	800,000	50,400	63.00	76.34	
TSUBAME2006 (w/360CSs)	11,200	810,000	79,430	98.06	72.32	
TSUBAME2007 (w/648CSs)	11,776	820,000	102,200	124.63	69.63	1.00
Earth Simulator	5120	6,000,000	40,000	6.67	1171.88	0.05
ASCI Purple (LLNL)	12240	6,000,000	77,824	12.97	490.20	0.10
AIST Supercluster (Opteron)	3188	522,240	14400	27.57	163.81	0.22
LLNL BG/L (rack)	2048	25,000	5734.4	229.38	12.21	1.84
Next Gen BG/P (rack)	4096	30,000	16384	546.13	7.32	4.38
TSUBAME 2.0 (2010Q3/4)	160,000	810,000	1,024,000	1264.20	5.06	10.14

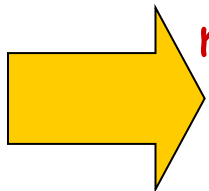
TSUBAME 2.0 x24 improvement in 4.5 years...? → ~ x1000 over 10 years

Scaling Peta to Exa Design?

- Shorten latency as much as possible
 - Extreme multi-core incl. vectors
 - "Fat" nodes, exploit short-distance interconnection
 - Direct cross-node DMA (e.g., put/get for PGAS)
- Hide latency if cannot be shortened
 - Dynamic multithreading (Old: dataflow, New: GPUs)
 - Trade Bandwidth for Latency (so we do need BW...)
 - Departure from simple mesh system scaling
- Change Latency-Starved Algorithms
 - From implicit Methods to direct/hybrid methods
 - Structural locality, extrapolation, stochastics (MC)
 - Need good programming model/lang/system for this...

GPUs as Commodity Massively Parallel Vector Processors

- E.g., NVIDIA Tesla, AMD Firestream
 - High Peak Performance > 1TFlops
 - Good for tightly coupled code e.g. Nbody
 - High Memory bandwidth (>100GB/s)
 - Good for sparse codes e.g. CFD
 - Low latency over shared memory
 - Thousands threads hide latency w/zero overhead
 - Slow and Parallel and Efficient vector engines for HPC
 - Restrictions: Limited non-stream memory access, PCI-express overhead, programming model etc.



How do we exploit them given vector computing experiences?

GPUs as Modern-Day Vector Engines

Two types of HPC Workloads

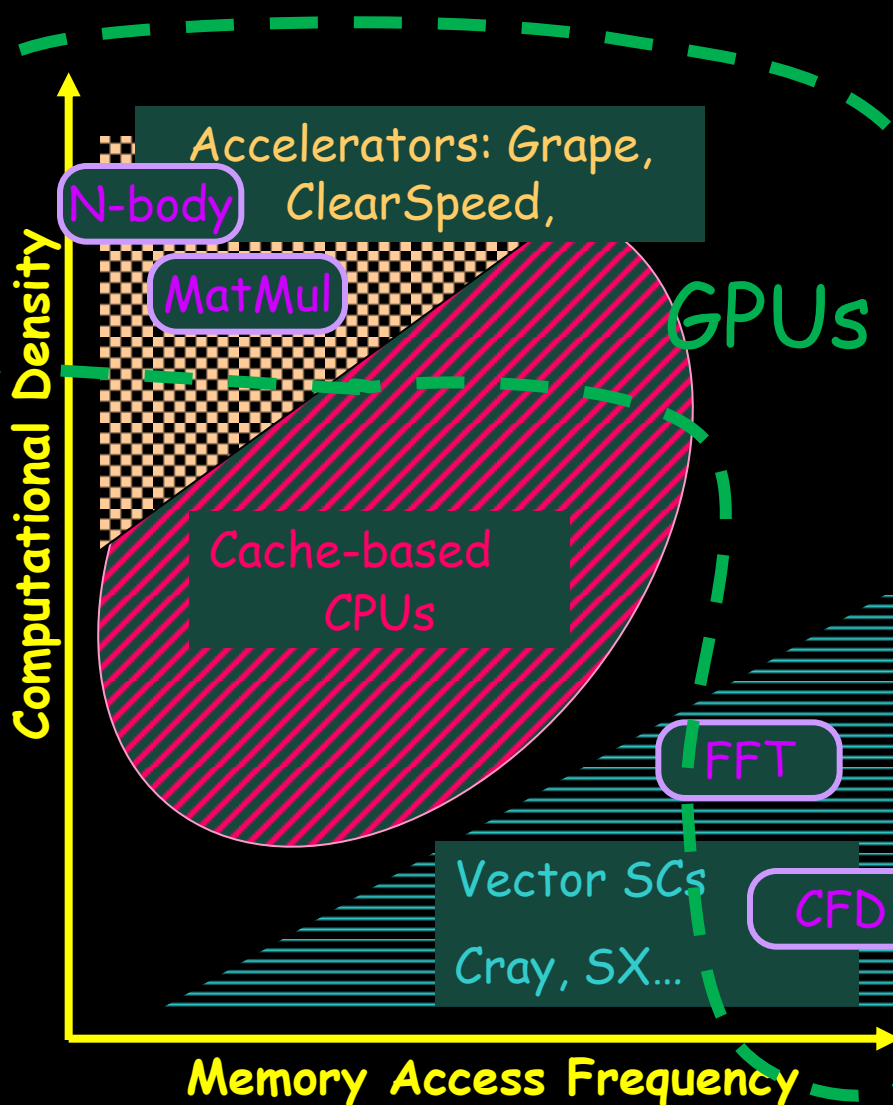
- High Computational Density
=> Traditional Accelerators
- High Memory Access Frequency
=> Traditional Vector SCs

Scalar CPUs are so-so at both
=> Massive system for speedup

GPUs are both modern-day vector engines and high compute density accelerators

=> Efficient Element of Next-Gen Supercomputers

Small memory, limited CPU-GPU BW, high-overhead communication with CPUs and other GPUs, lack of system-level SW support incl. fault tolerance, programming, ...



Our Research@
Tokyo Tech

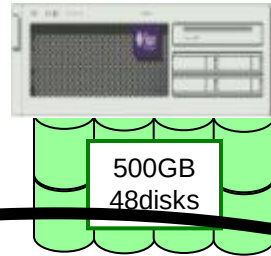
TSUBAME 1.2 Experimental Evolution (Oct. 2008)

The first "Petascale" SC in Japan



Voltaire ISR9288 Infiniband x8
10Gbps x2 ~1310+50 Ports
~13.5Terabits/s
(3Tbits bisection)

NEC SX-8i



Storage

1.5 Petabyte (Sun x4500 x 60)

0.1Petabyte (NEC iStore)

Lustre FS, NFS, CIFS, WebDAV (over IP)

60GB/s aggregate I/O BW

Sun x4600 (16 Opteron Cores)

32~128 GBytes/Node

10480core/655Nodes

21.4TeraBytes

50.4TeraFlops

OS Linux (SuSE 9, 10)

NAREGI Grid MW

**87.1TF Linpack (The First GPU
Supercomputer on Top500)
[IPDPS10]**

NEW Deploy:
GCOE TSUBASA
Harpertown-Xeon
90Node 720CPU
8.2TeraFlops

PCI-e

ClearSpeed CSX600

SIMD accelerator

648 boards,

52.2TeraFlops

SFP/DFP

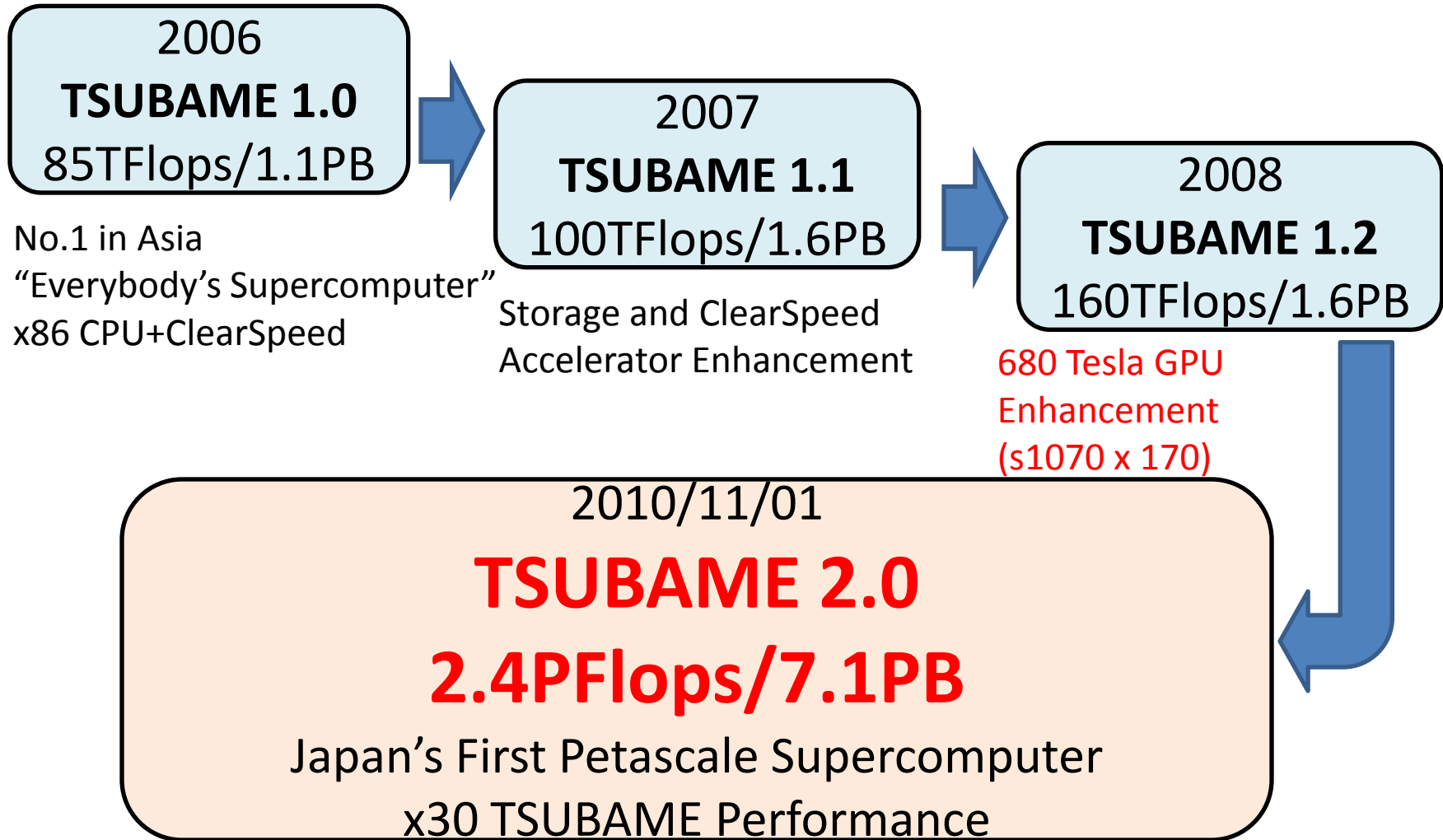
170 Nvidia Tesla 1070, ~680 Tesla cards
High Performance in Many BW-Intensive Apps
10% power increase over TSUBAME 1.0



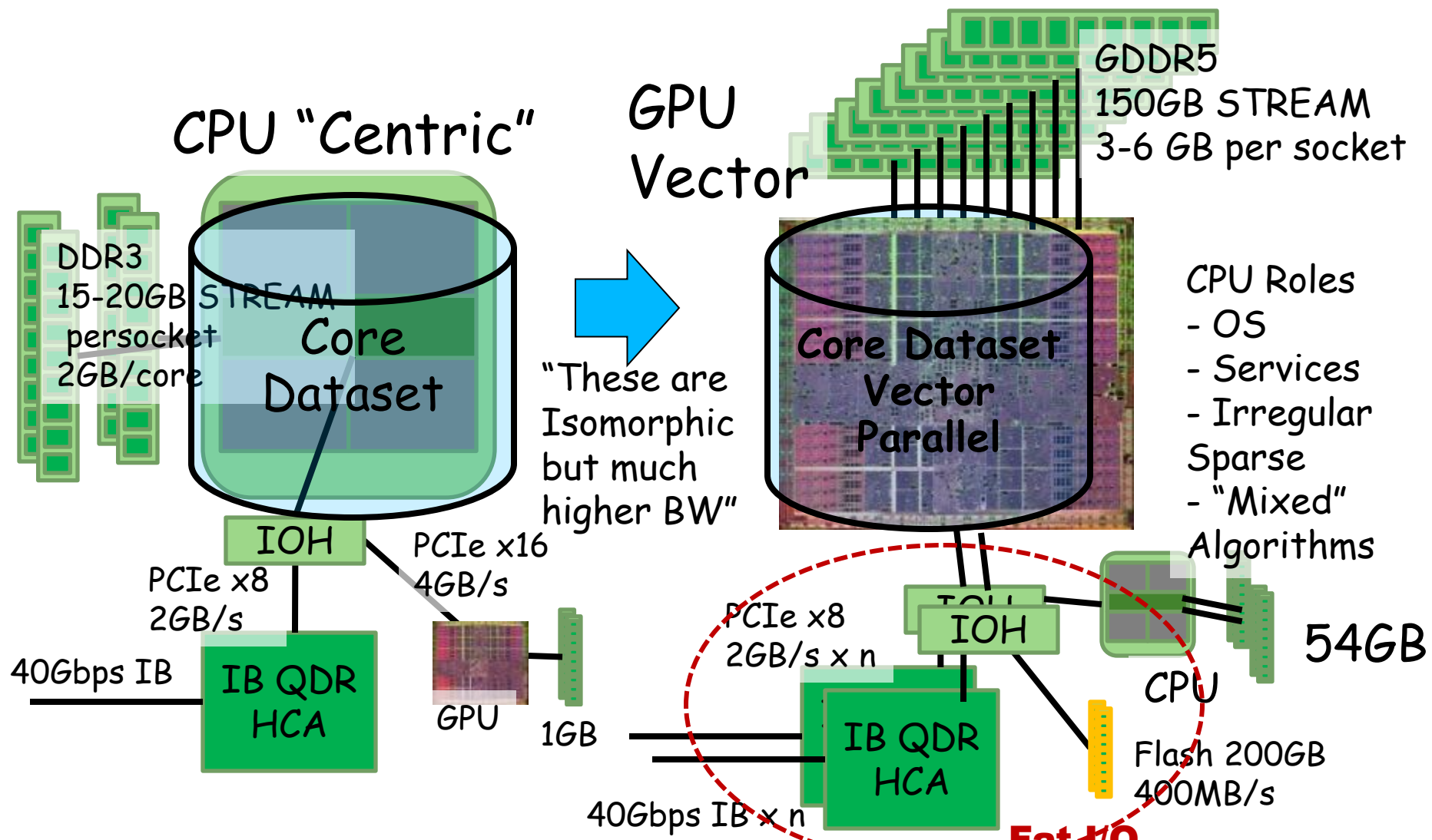
680 Unit Tesla Installation...
While TSUBAME in Production Service (!)



History of TSUBAME



From TSUBAME 1.2 to 2.0: From CPU Centric to GPU Centric Nodes for Scaling High Bandwidth in Network & I/O!!



Tsubame 2: Start with Clean Slate!!!

- Achieve 30x performance over Tsubame 1.0
- Support diverse research workloads
 - Weather, informatics, protein modeling, etc.
- Balanced design for scale
 - Bandwidth to GPUs, CPU and GPU high memory capacity
 - High-bandwidth, low latency, full bi-section network
 - Solid state disks for high bandwidth, scalable storage I/O
- Green: stay within power threshold
 - System estimated to require 1.8 MW with PUE of 1.2
- Density (GPU per CPU footprint) for costs and space
- Expand support for Tokyo Tech's 2000 users; provide dynamic "cloud-like" provisioning for both Microsoft® Windows® and Linux workloads

TSUBAME2.0 Global Partnership

NEC: Main Contractor, overall system integration, Cloud-SC mgmt.

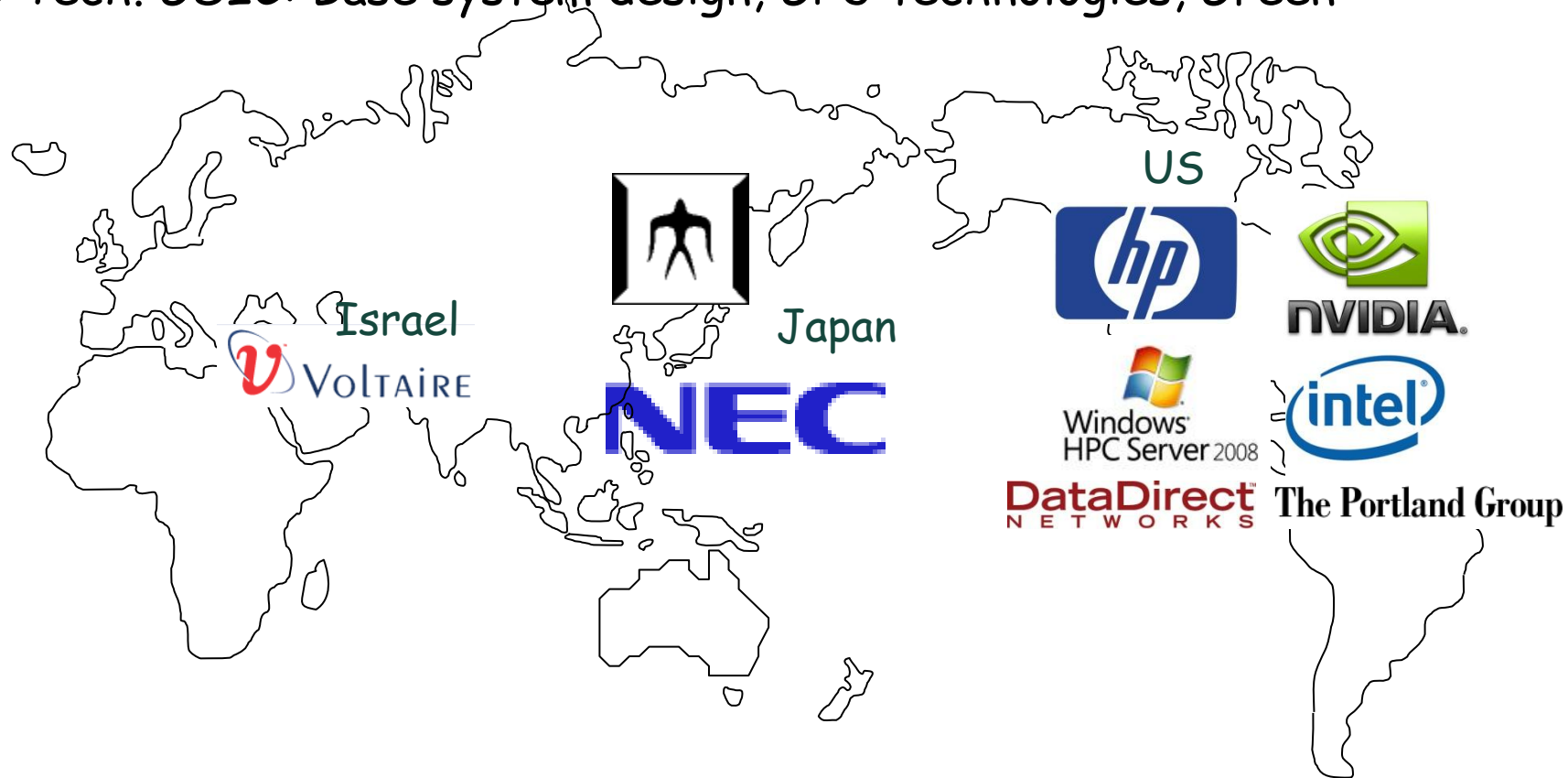
HP: Node R&D, Green; Microsoft: WindowsHPC, Cloud & VM

NVIDIA: Fermi GPUs, CUDA; Voltaire: QDR Infiniband Network

DDN: Large-Scale Storage; Intel: Westmere & Nehalem-EX CPUs

PGI: GPU Vectorizing Compiler

Tokyo Tech. GSIC: Base system design, GPU technologies, Green



Highlights of TSUBAME 2.0 Design (Oct. 2010) w/NEC-HP

- 2.4 PF multi-core x86 + Fermi GPGPU
 - ▶ 1432 nodes, Intel Westmere/Nehalem EX
 - ▶ 4224 NVIDIA Tesla (Fermi) M2050 GPUs
 - ▶ ~80,000 total CPU and GPU "cores", High Bandwidth
 - ▶ 1.9 million "CUDA cores", $32K \times 4K = 130$ million CUDA threads(!)
- 0.72 Petabyte/s aggregate mem BW,
 - ▶ Effective 0.3-0.5 Bytes/Flop, restrained memory capacity (100TB)
- Optical Dual-Rail IB-QDR BW, full bisection BW(Fat Tree)
 - ▶ 200Tbits/s, Likely fastest in the world, still scalable
- Flash/node, ~200TB (1PB in future), 660GB/s I/O BW
 - ▶ >7 PB IB attached HDDs, 15PB Total HFS incl. LTO tape
- Low power & efficient cooling, comparable to TSUBAME 1.0 (~1MW); PUE = 1.28 (60% better c.f. TSUBAME1)
- Virtualization and Dynamic Provisioning of Windows HPC + Linux job migration etc



TSUBAME2.0 Nov. 1, 2011

See Live @ Tokyo Tech Booth #1127



TSUBAME2.0: A GPU-centric Green 2.4 Petaflops Supercomputer

Tsubame 2.0: "Tiny" footprint, very power efficient

- Floorspace less than 200m² (2,100 ft²)
- Top-class power efficient machine on the Green 500

System

(42 Racks)

1408 GPU Compute Nodes,

34 Nehalem "Fat Memory" Nodes

Rack

(8 Node Chassis)



2.4 PFLOPS
80 TB

Node Chassis

(4 Compute Nodes)



6.7 TFLOPS
220 GB/412 GB

Compute Node

(2 CPUs, 3 GPUs)



1.6 TFLOPS
55 GB/103 GB

Chip

(CPU, GPU)



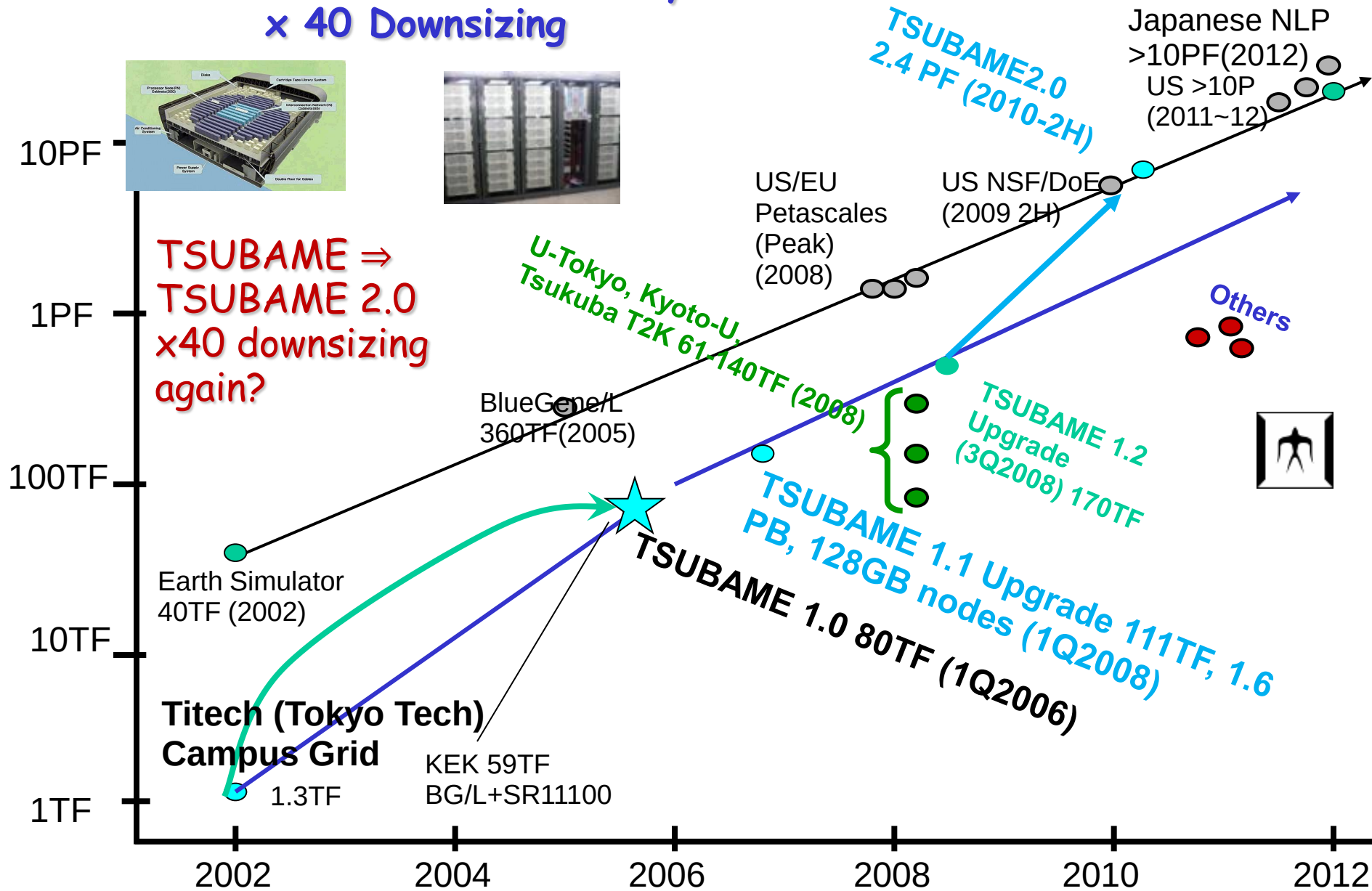
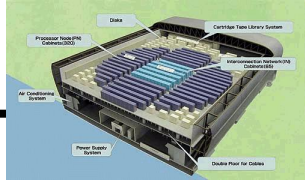
CPU(Westmere EP)
76.8 GFLOPS

GPUs(Tesla M2050)
515 GFLOPS
3 GB

Integrated by NEC Corporation

TSUBAME 2.0 Performance

Earth Simulator $\Rightarrow$ TSUBAME 4years
 $\times$ 40 Downsizing



TSUBAME2.0 World Rankings

(Nov. 2010 PREDICTION for Green500!!!)

The Top 500 (Absolute Performance)

- #1 : ~2.5 PetaFlops: China Defense Univ. Dawning Tianhe 1-A
- #2 : 1.76 Petaflops: US ORNL Cray XT5 Jaguar
- #3 : 1.27 PetaFlops: China Shenzen SC Nebulae
- #4 : 1.19 PetaFlops: Japan Tokyo Tech. HP/NEC TSUBAME2.0
- #5 : 1.054 PetaFlops: US LLBL Cray XE6 Hopper
- #~33 (#2 Japan) : 0.191 Petaflops : JAEA Fujitsu



Top500 BoF
Tue 16th
17:30~

The Green 500 (Performance/Power Efficiency)

- #1: ???
- #2: ???
- #3: ???
- #4: ???
- #5: ???

TSUBAME2.0
> 900MFlops/W



Green 500
BoF
Thu 18th
12:15~

“Little Green 500” collaboration w/Microsoft
Achieved 1.037 Gigaflops/W in power efficient experimental config.

Petaflops? Gigaflops/W?



**x66,000
faster
x3 power
efficient**



x44,000 Data



Laptop: SONY Vaio type Z (VPCZ1)
CPU: Intel Core i7 620M (2.66GHz)
MEMORY: DDR3-1066 4GBx2
OS: Microsoft Windows 7 Ultimate 64bit
HPL: Intel(R) Optimized LINPACK Benchmark
for Windows (10.2.6.015)
256GB HDD

18.1 GigaFlops Linpack
369 MegaFlops/W

Supercomputer: TSUBAME 2.0
CPU: 2714 Intel Westmere 2.93 Ghz
GPU: 4071 nVidia Fermi M2050
MEMORY: DDR3-1333 80TB + GDDR5 12TB
OS: SuSE Linux 11 + Windows HPC Server
R2

HPL: Tokyo Tech Heterogeneous HPL
11PB Hierarchical Storage
1.192 PetaFlops Linpack
1043 MegaFlops/W

TSUBAME2.0 System Overview (2.4 Pflops/15PB)

Petascale Storage : Total **7.13PB**(Lustre + Accelerated NFS Home)

Lustre Partition 5.93PB

x5

MDS,OSS
HP DL360 G6 30nodes
Storage
DDN SFA10000x5
(10 enclosures x5)
Lustre (5 Filesystems)
OSS: 20 OST: 5.9PB
MDS: 10 MDT: 30TB

OSS x20 MDS x10

**Home NFS/
iSCSI**

Storage Server
HP DL380 G6 4nodes
BlueArc Mercury 100 x2
Storage
DDN SFA10000 x1
(10 enclosures x1)

NFS,CIFS x4 NFS,CIFS,iSCSI
Accelerationx2

**Tape System
Sun SL8500 8PB**

SuperTitenet

SuperSinet3

Node Interconnect: Optical, Full Bisection, Non-Blocking, Dual-Rail QDR Infiniband

Core Switch

12 switches

Edge Switch

179 switches

Edge Switch(w/10GbE)

6 switches

Voltaire Grid Director 4700
IB QDR: 324ports

Voltaire Grid Director 4036
IB QDR : 36 ports

Voltaire Grid Director 4036E
IB QDR:34ports
10GbE: 2ports

Mgmt Servers

Compute Nodes : 2.4PFlops(CPU+GPU) 224.69TFlops(CPU)

"Thin" Nodes

NEW DESIGN Hewlett Packard CPU+GPU
High BW Compute Nodes x 1408
Intel Westmere-EP 2.93GHz
(TB 3.196GHz) 12Cores/node
Mem:55.8GB(=52GiB)or 103GB(=96GiB)
GPU NVIDIA M2050 515GFlops,
3GPUs/node
SSD 60GB x 2 120GB ※55.8GB node
120GB x 2 240GB ※103GB node
OS: Suse Linux Enterprise + Windows

HPC
4224 NVIDIA "Fermi" GPUs
Memory Total : 80.55TB
SSD Total : 173.88TB

1408nodes
(32node x44 Racks)

"Medium" Nodes

HP 4Socket Server 24nodes
CPU Nehalem-EX 2.0GHz
32Cores/node
Mem:137GB(=128GiB)
SSD 120GB x 4 480GB
OS: Suse Linux Enterprise

24 nodes 6.14TFLOPS

"Fat" Nodes

HP 4Socket Server 10nodes
CPU Nehalem-EX 2.0GHz
32Core/node
Mem:274GB(=256GiB)x8
549GB(=512GiB) x2
SSD 120GB x 4 480GB
OS: Suse Linux Enterprise

10 nodes 2.56TFLOPS

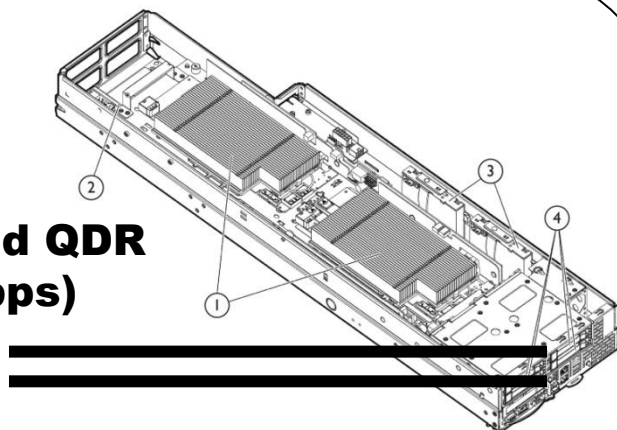
PCI-E gen2 x16 x2slot/node

GSIC:NVIDIA Tesla S1070GPU (34 units)

TSUBAME2.0 Compute Nodes

Thin Node

**Infiniband QDR
x2 (80Gbps)**



**HP SL390G7 (Developed for
TSUBAME 2.0)**

**GPU: NVIDIA Fermi M2050 x 3
515GFlops, 3GByte memory /GPU
CPU: Intel Westmere-EP 2.93GHz x2
(12cores/node)**

Memory: 54, 96 GB DDR3-1333

SSD : 60GBx2, 120GBx2

IB QDR



**PCI-e Gen2x16 x2
NVIDIA Tesla
S1070 GPU**

HP 4 Socket Server

**CPU: Intel Nehalem-EX 2.0GHz x4
(32cores/node)**

Memory: 128, 256, 512GB DDR3-1066

SSD : 120GB x4 (480GB/node)

1408nodes :

**4224GPUs : 59,136 SIMD Vector
Cores, 2175.36TFlops (Double FP)**

**2816CPUs, 16,896 Scalar Cores:
215.99TFlops**

Total : 2391.35TFLOPS

**Memory : 80.6TB (CPU) + 12.7TB
(GPU)**

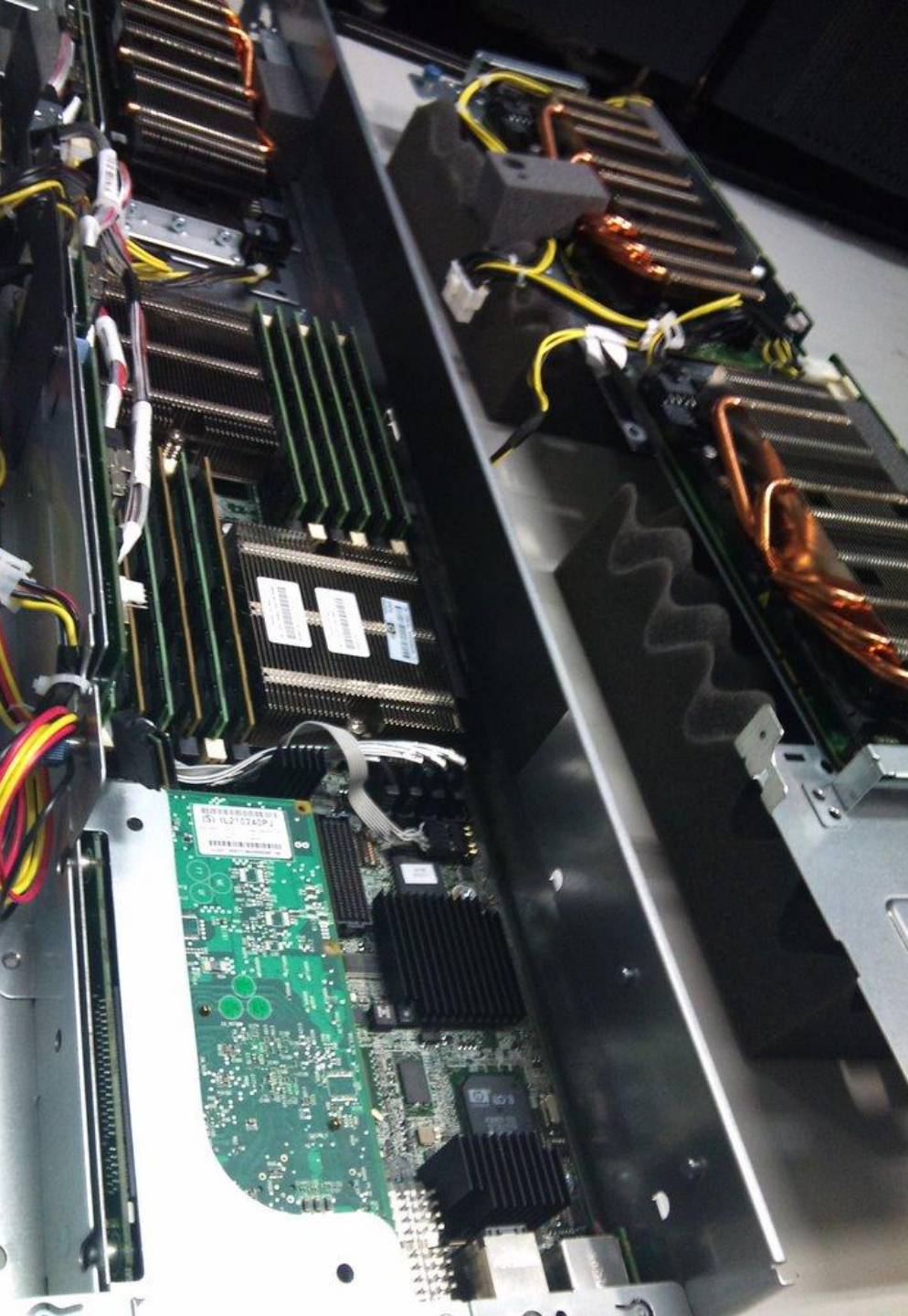
SSD : 173.9TB

**34 nodes:
8.7TFlops**

**Memory:
6.0TB+GPU**

SSD: 16TB+

**Total Perf
2.4PFlops
Mem : ~100TB
SSD: ~200TB**



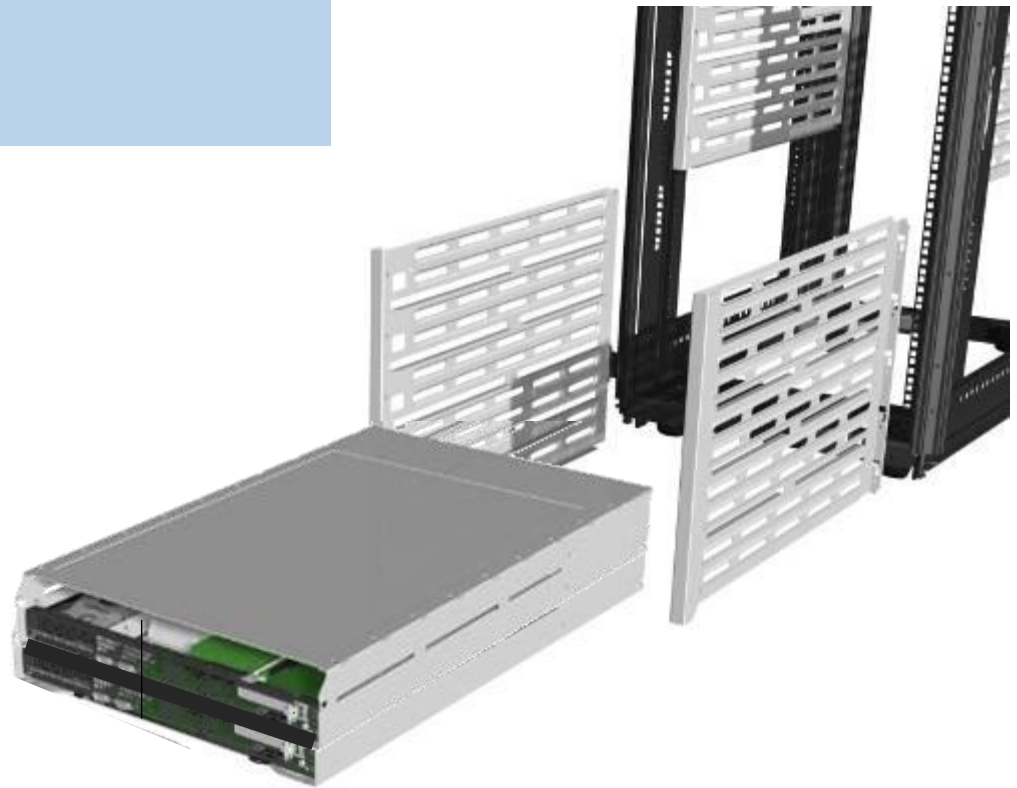


GPU-optimized node in HP Skinless Packaging

Typical Tsubame2 node:

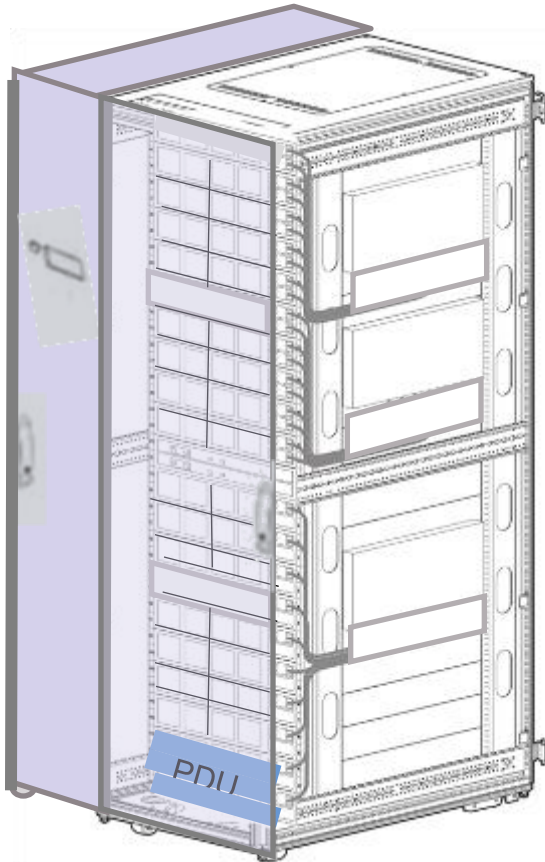
- 2 Intel® Xeon® Processor X5670 (six-core 2.93Ghz)
- 54 GB or 96 GB memory
- 2 SSDs per node
- Dual Rail IB
- 3 NVIDIA M2050 GPU board

- Nodes go into chassis with shared fans and power
- Chassis mounted into lightweight racks
- Easy assembly
- Lower weight
- Reduced heat retention
- Modular and flexible
- Standard 19" racks



Compute Rack Building Block for Tsubame2

HP Modular Cooling System
G2 rack



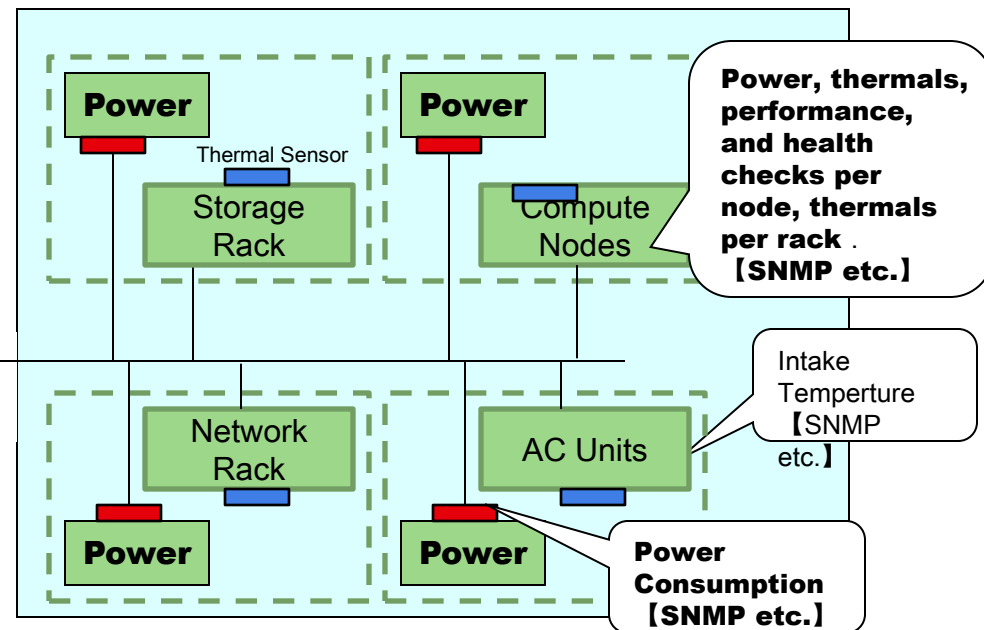
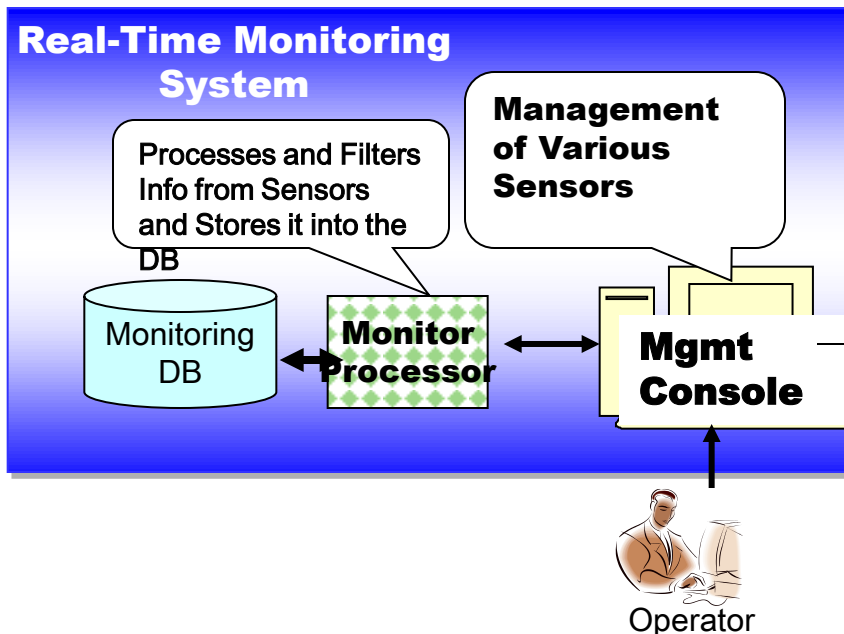
- 42U HP Modular Cooling System G2 rack
 - 30 nodes per rack – 90 GPUs
 - 8 chassis with Advanced Power Management
 - 1 HP Network Switch for shared console and admin local area network
 - 2 Airflow Dam (Liquid Cooled)
 - 4 Voltaire 4036 Leaf Switch
 - Power distribution units
- Power per rack approximately

35KW

Green and Reliable SC

- ▶ Monitoring Sensors (incl. GPUs)
 - ▶ Thermals
 - ▶ Power Consumption
 - ▶ Utilization
 - ▶ HW and SW Health checks

- **Advanced Power Management**
 - **Dynamic Power Capping**
 - **Power monitoring**
 - **Node level power off/on**
- **Thermal logic technologies**
 - **Shared power and fans**
 - **Energy-efficient fans**
 - **Three-phase load balancing**
- **94% Platinum Common Slot Power Supplies**



Cooling: Enhanced HP Modular Cooling System G2

HP's **water-cooled** rack

Completely closed racks with their own heat exchanger.

1.5 x width of normal rack+rear ext.

Cooling for high density deployments

35kW of cooling capacity single rack

- **Highest Rack Heat Density ever**
- **3000CFM Intake airflow with 7C chiller water**

up to 2000 lbs of IT equipment

Uniform air flow across the front of the servers

Automatic door opening mechanism controlling both racks

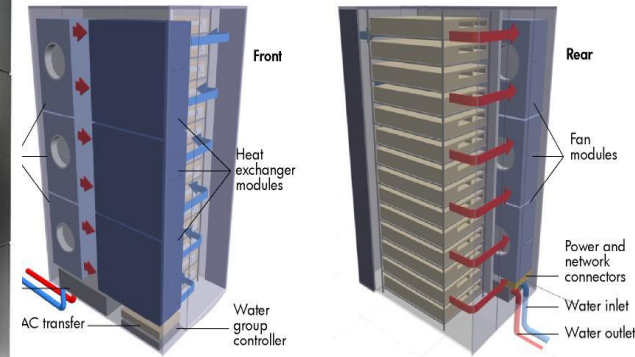
Adjustable temperature set point

Removes 95% to 97% of heat inside racks

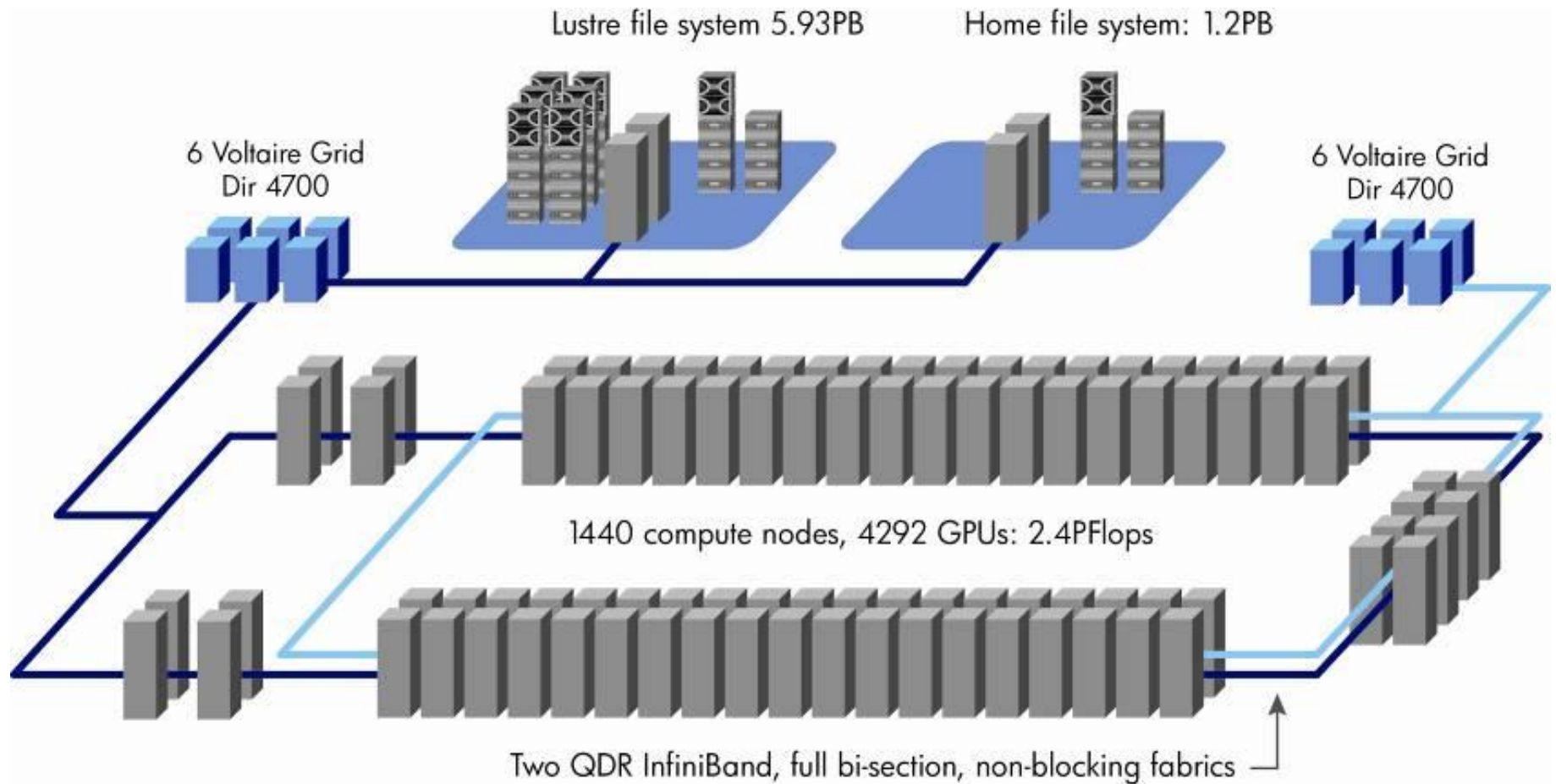
Polycarbonate front door reduces ambient noise considerably



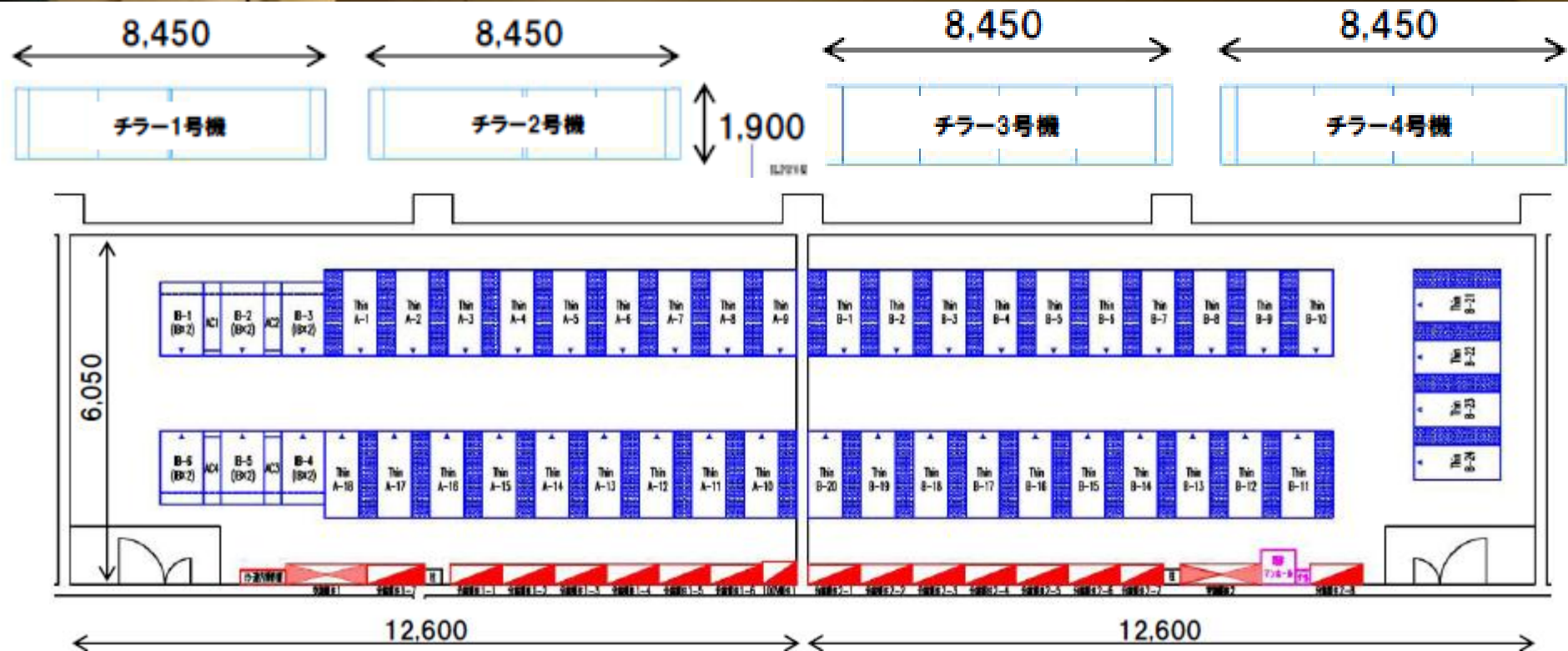
~= Entire Earth Simulator (rack = 50TF)



Pulling It All Together



TSUBAME2.0 Layout (200m2 for main compute nodes)







**2.4 Petaflops, 1408 nodes
~50 compute racks + 6 switch racks**

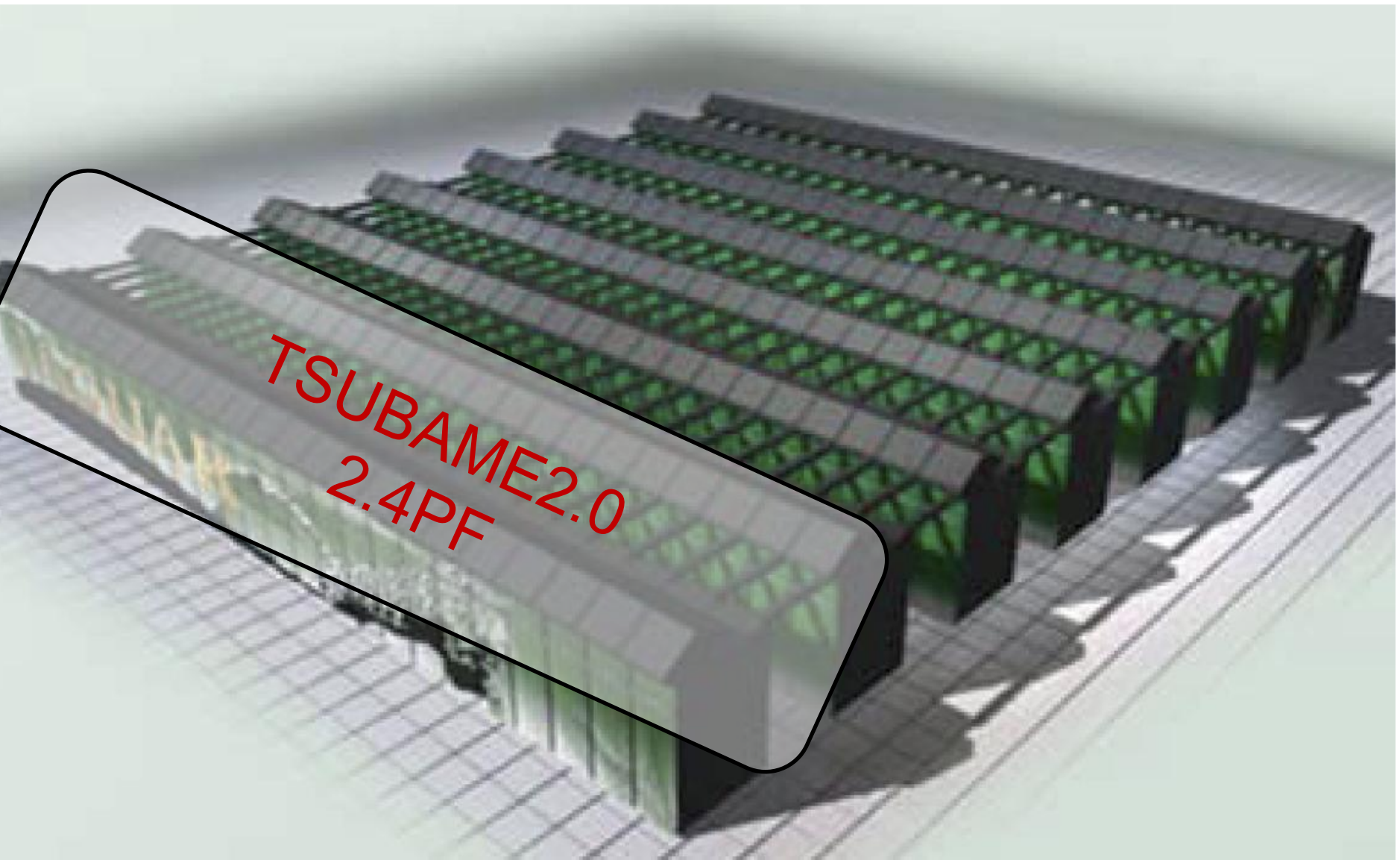


**Two Rooms, Total 160m²
1.4MW (Max, Linpack), 0.48MW (Idle)**

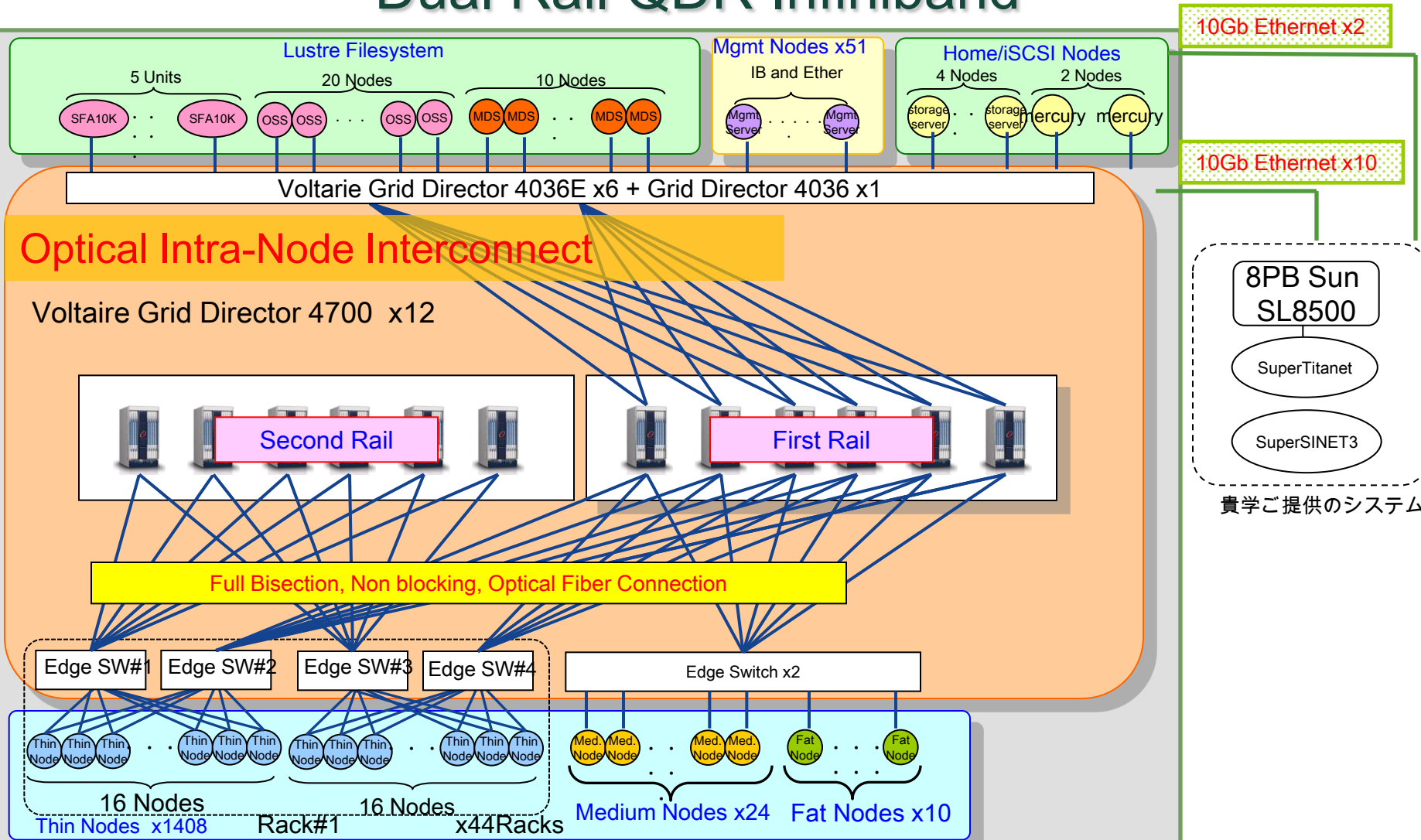


ORNL Jaguar and Tsubame 2.0

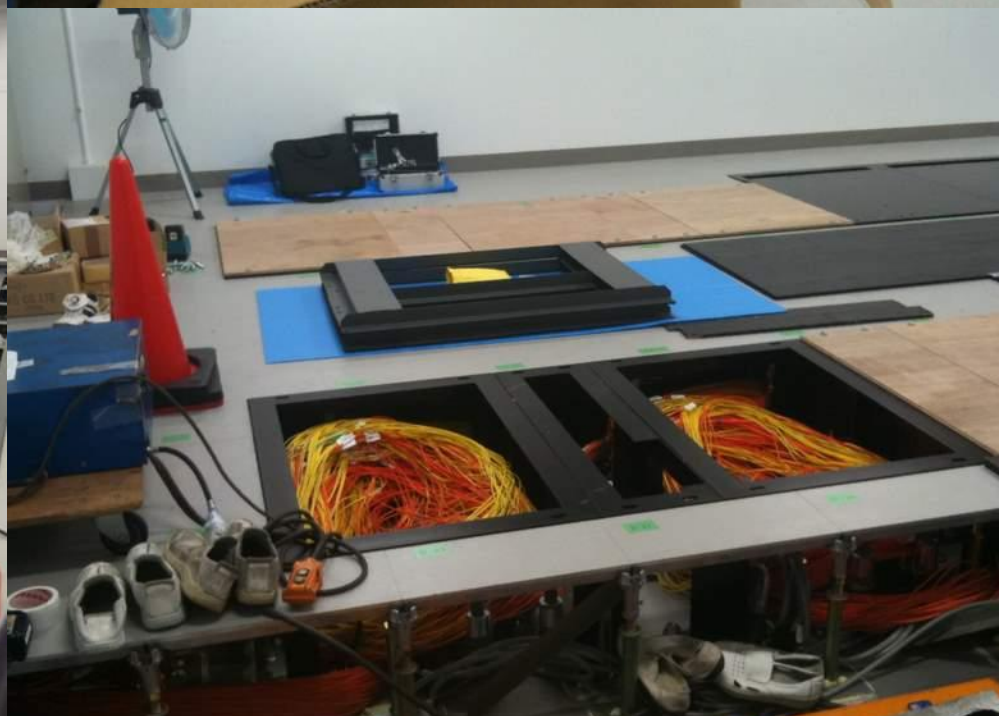
Similar Peak Performance, 1/4 the Size and Power



TSUBAME 2.0 Full Bisection Fat Tree, Optical, Dual Rail QDR Infiniband



TSUBAME2.0ネットワーク全体図



TSUBAME2.0 Storage

Multi-Petabyte storage consisting of Lustre Parallel Filesystem Partition
and NFS/CIFS/iSCSI Home Partition + Node SSD Acceleration

Lustre Parallel Filesystem Partition,

MDS : HP DL360 G6 **5.96PB**

- CPU: Intel Westmere-EP x2 socket (12 Cores)
- Memory : 51GB(=48GiB)
- IB HCA: IB 4X QDR PCI-e G2 x1 port

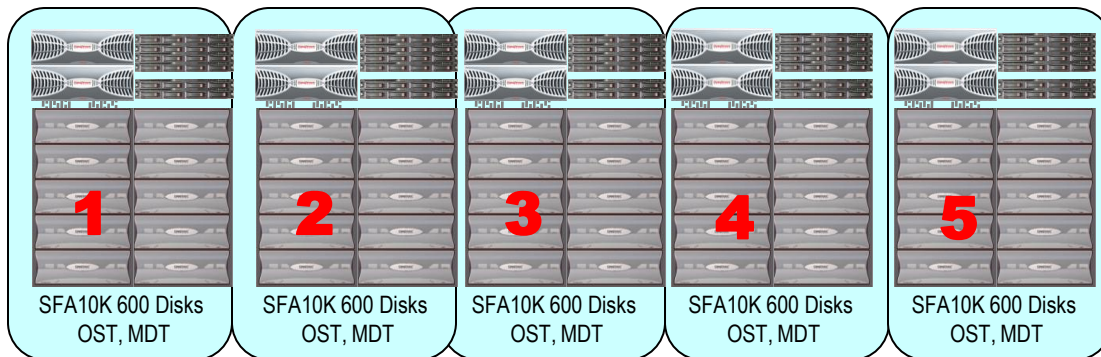
OSS : HP DL360 G6 x20

- CPU: Intel Westmere-EP x2 socket (12 Cores)
- Memory : 25GB(=24GiB)
- IB HCA : IB 4X QDR PCI-e G2 x2 port

Storage : DDN SFA10000 x5

- Total Capacity : 5.93PB

2TB SATA x 2950 Disks + 600GB SAS x 50 Disks



Home Partition 1.2PB

NFS/CIFS : HP DL380 G6 x4

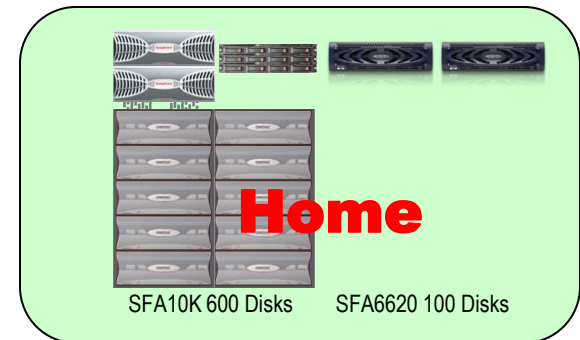
- CPU: Intel Westmere-EP x2 socket (12 Cores)
- Memory : 51GB(=48GiB)
- IB HCA: IB 4X QDR PCI-e G2 x2 port

NFS/CIFS/iSCSI : BlueArc Mercury100 x2

- 10GbE x2

ストレージ : DDN SFA10000 x1

- Total Capacity : 1.2PB
- 2TB SATA x 600 Disks



7.13PB HDD + 200TB SSD + 8PB Tape

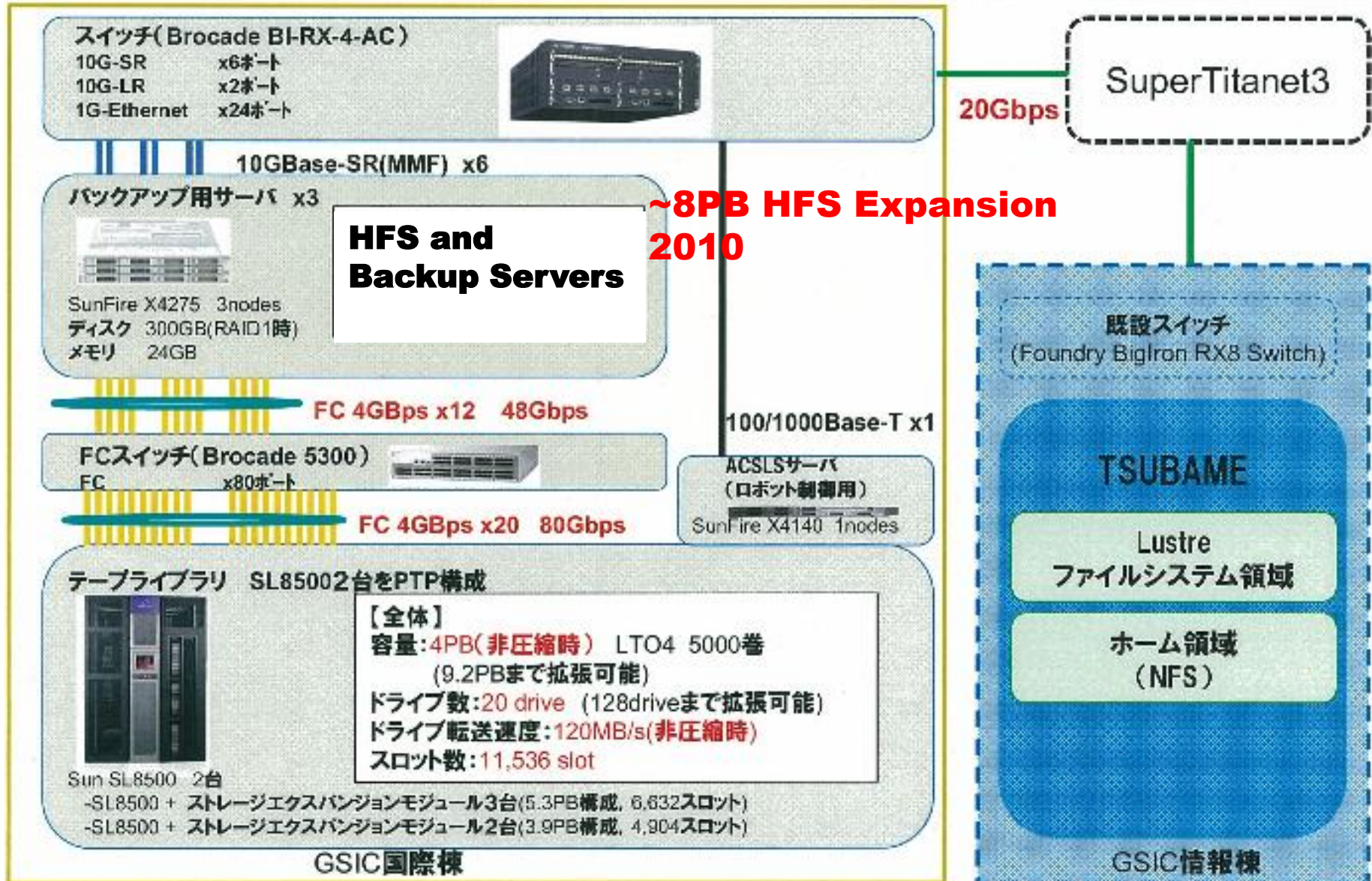


200 TB SSD (SLC, 1PB MLC future)
7.1 PB HDD (Highly redundant)
4-8 PB Tape (HFS+Backup)
All Infiniband+10GbE Connected

Lustre + GPFS
NFS, CIFS, WebDAV,...
Tivoli HFS
GridScaler + BlueArc

TSUBAME2.0 Tape System

HFS of over15PB and Scalable



RENKEI-POP Distributed Storage

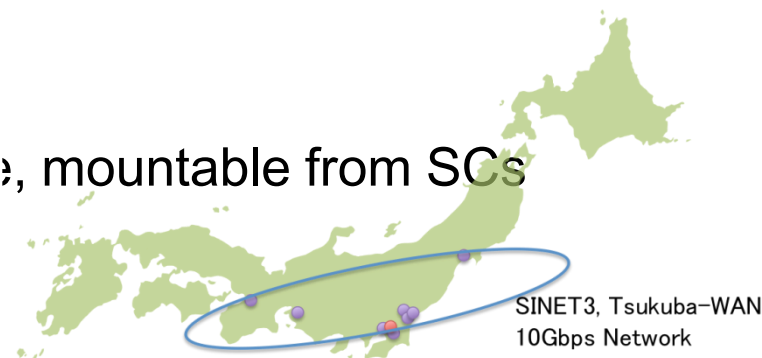
- Objective: SC Centers-Wide Distributed Storage on SINET National NW
 - ▶ RENKEI (MEXT e-Science) Project and R&D and nationwide deployment of RENKEI-PoPs (Point of Presence)
 - ▶ Data transfer server engine with large capacity and fast I/O
 - ▶ “RENKEI-Cloud” on SINET3 with various Grid/Cloud Services including Gfarm distributed filesystem



CPU	Core i7 975 Extreme (3.33 GHz)
Memory	12GB (DDR3 PC3-10600 , 2GB*6)
NIC	10GbE (without TCP/IP Offload Engine)
System Disk	500GB HDD
SSD RAID	30TB (RAID 5, 2TB HDD x 16)

Tokyo Inst. Tech.	Osaka Univ.
National Inst. Inform.	KEK
Nagoya Univ.	Univ. Tsukuba
AIST	Tohoku Univ.

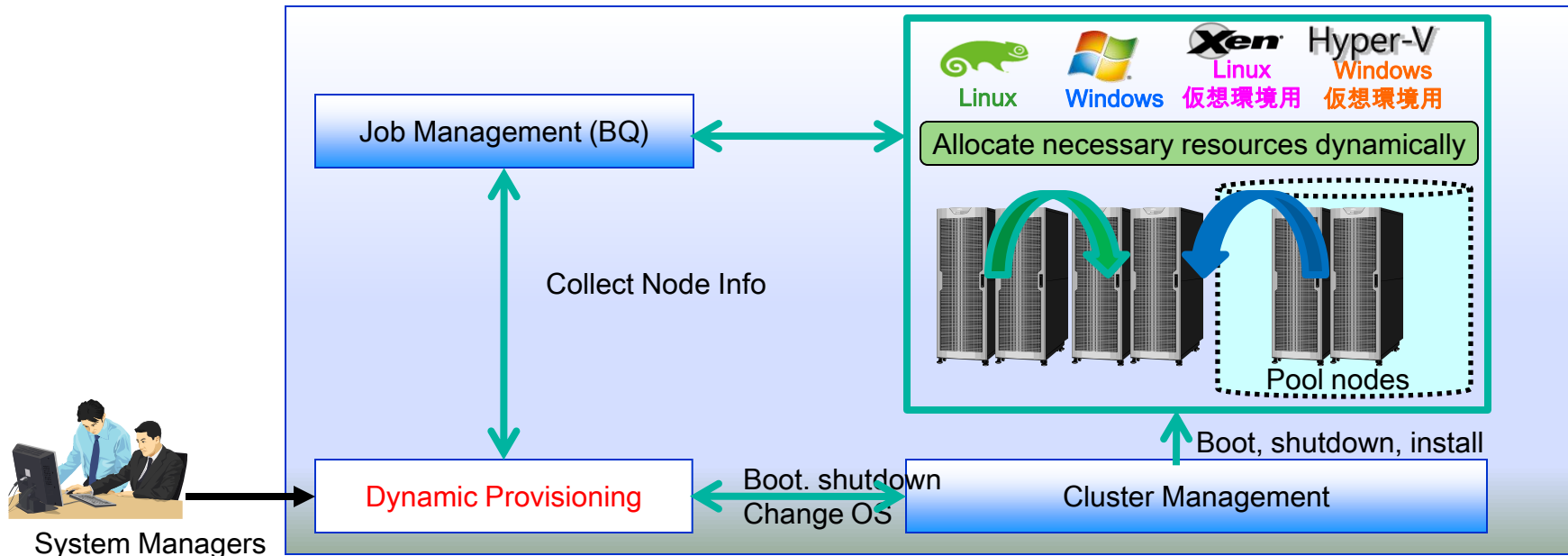
, mountable from SCs



Cloud Provisioning Supercomputing Nodes

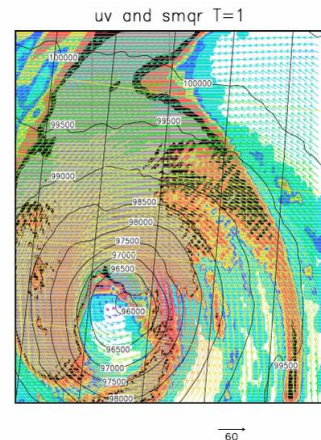
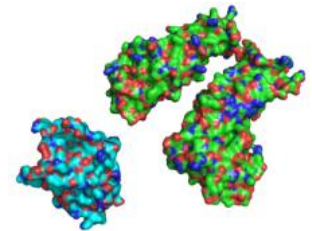
Tsubame 2.0 will have a dynamic provisioning system that works with batch queue systems and cluster management systems to dynamically provision compute resources

- ▶ Pool nodes can be dynamically allocated to various supercomputing services
- ▶ Linux and Windows batch queue systems coordinate to manage the nodes
- ▶ Dynamic increase/decrease of nodes allocated to particular services
 - ※ Increase / Decrease of Windows HPC and Linux nodes of different configurations
- ▶ Virtual Nodes on VMs on respective OSes can be subject to dynamically allocated and scheduled.

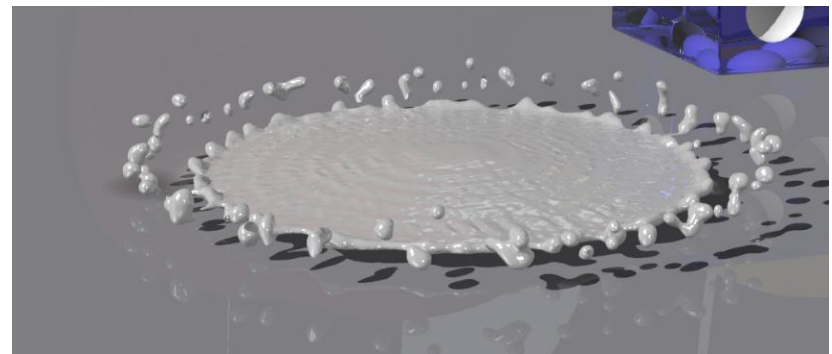
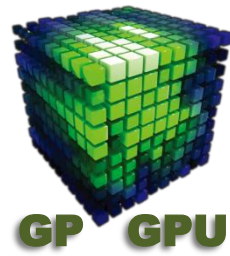


TSUBAME2.0 Estimated Performances

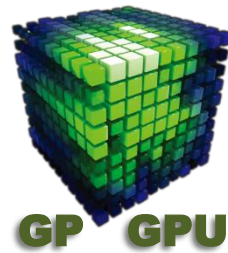
- 1.192 TFlops Linpack [IEEE IPDPS 2010]
 - Top ranks Green 500?
- ~0.5 PF 3D Protein Rigid Docking (Node 3-D FFT) [SC08, SC09]
- 145Tlops ASUCA Weather Forecast [SC10 Best Student Paper Finalist]
- Multiscale Simulation of Cardiovascular flows [SC10 Gordon Bell Finalist]
- Various FEM, Genomics, MD/MO (w/IMS), CS-Apps: search, optimization,



Gas-Liquid Two-Phase Flow



Fluid Equations



- Navier-Stokes solver : Fractional Step
- Time integration : 3rd TVD Runge-Kutta
- Advection term : 5th WENO
- Diffusion term : 4th FD
- Poisson : AMG-BiCGstab
- Surface tension : CSF model
- Surface capture : CLSVOF(THINC + Level-Set)

Continuum eq.

$$\nabla \cdot \mathbf{u} = 0$$

Momentum eq.

$$\frac{\partial \mathbf{u}}{\partial t} + (\mathbf{u} \cdot \nabla) \mathbf{u} = -\frac{1}{\rho} \nabla p + \nu \nabla^2 \mathbf{u} + \frac{1}{\rho} \mathbf{F}$$

Level-Set advection

$$\frac{\partial \phi}{\partial t} + (\mathbf{u} \cdot \nabla) \phi = 0$$

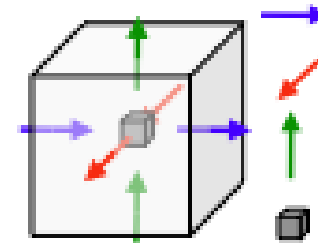
VOF continuum eq.

$$\frac{\partial \psi}{\partial t} + \nabla \cdot (\mathbf{u} \psi) = 0$$

Level-Set
re-initialization

$$\frac{\partial \phi}{\partial \tau} = \text{sgn}(\phi) (1 - |\nabla \phi|)$$

Staggered variable position



$\mathbf{u}$: velocity on x-direction

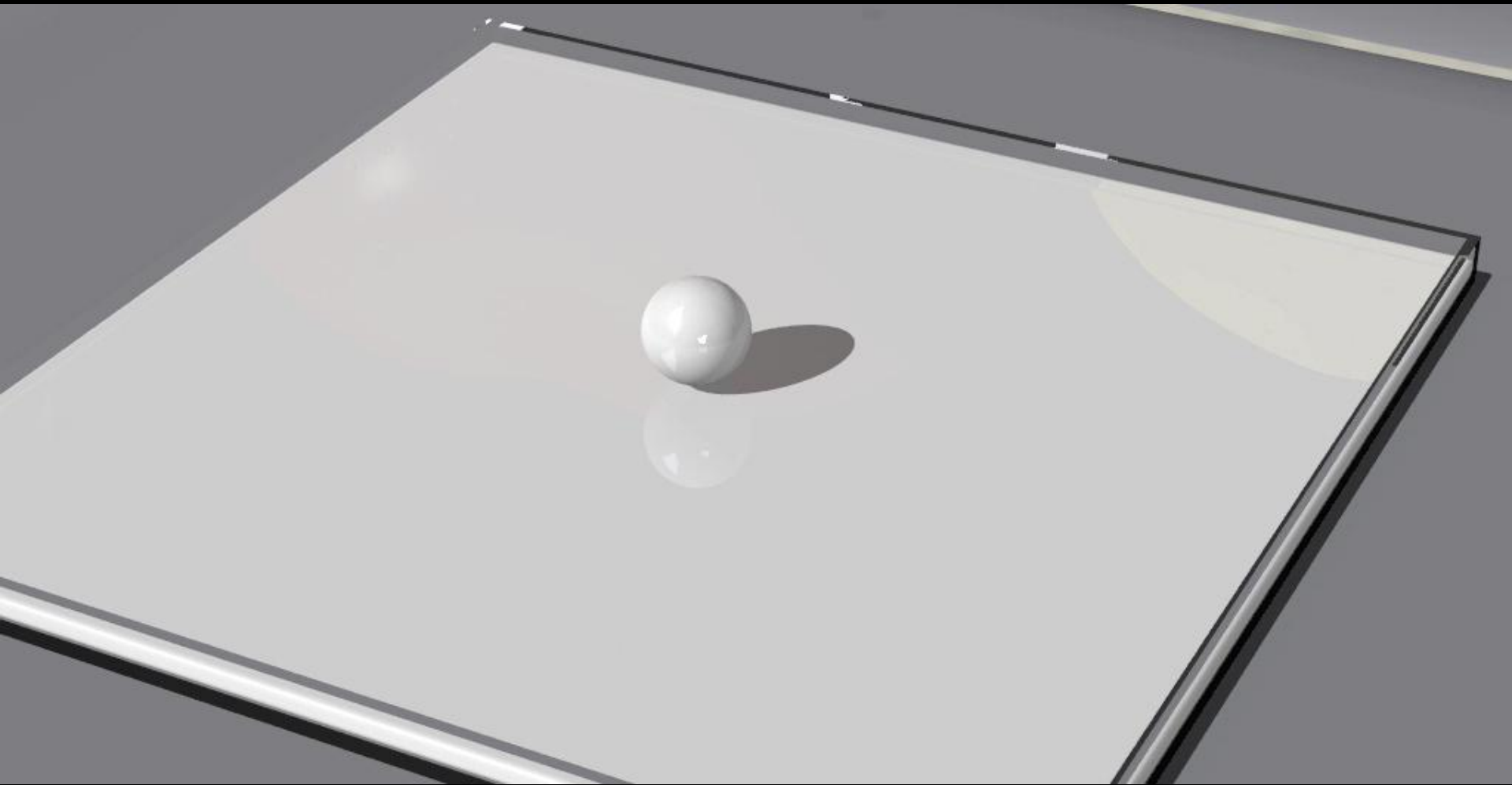
$\mathbf{v}$: velocity on y-direction

$\mathbf{w}$: velocity on z-direction

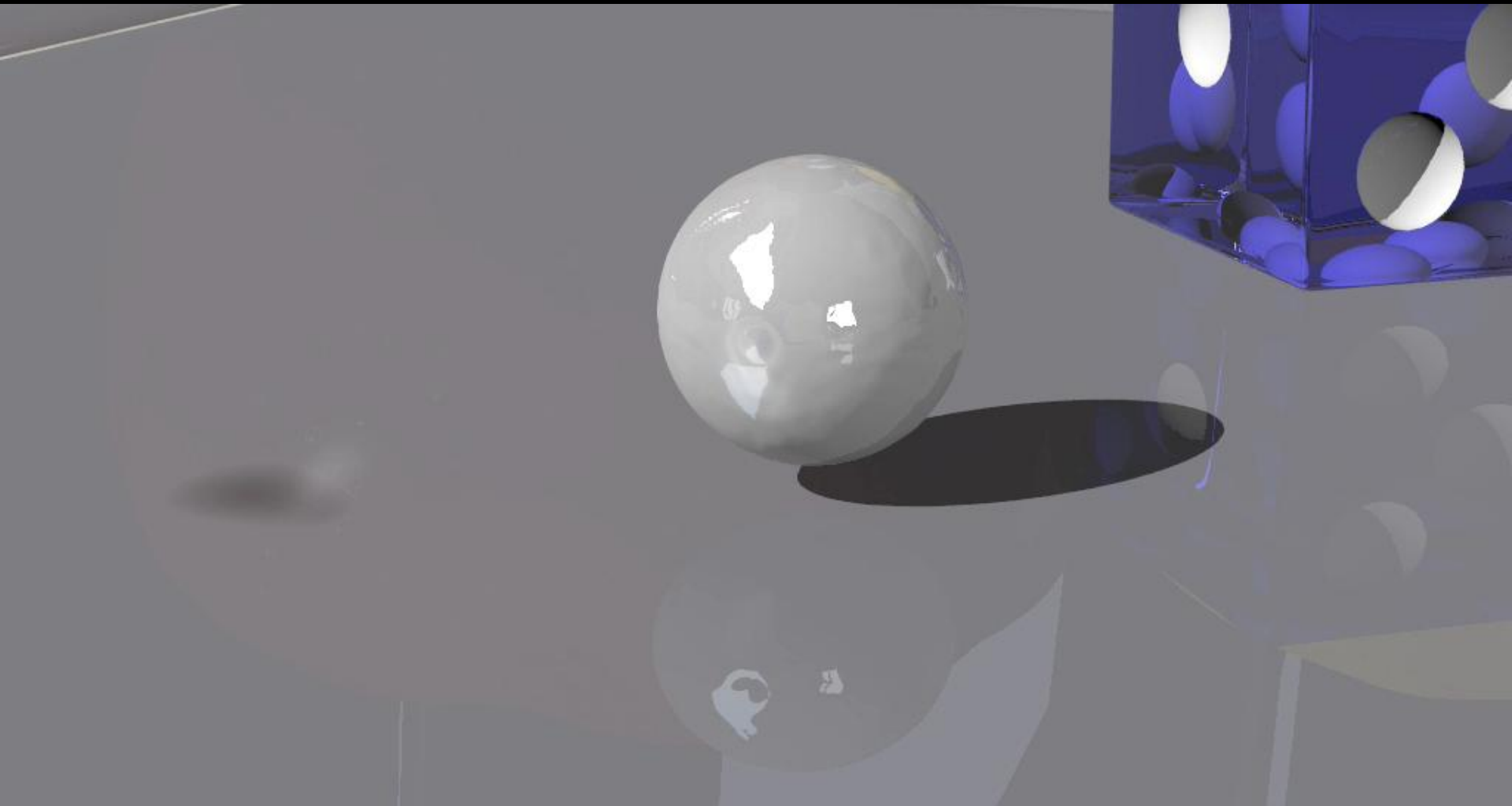
p, ϕ, ψ : scalar variables



Milk Crown



Drop on dry floor



Dendrite Solidification

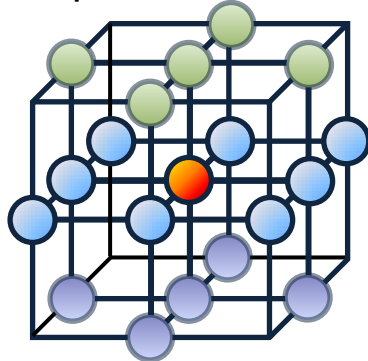
Allen-Cahn Equation based on Phase Field Model

$$\frac{\partial \phi}{\partial t} = M \left[\frac{\partial}{\partial x} \left(\varepsilon^2 \frac{\partial \phi}{\partial x} + \varepsilon \frac{\partial \varepsilon}{\partial \phi_x} |\nabla \phi|^2 \right) + \frac{\partial}{\partial y} \left(\varepsilon^2 \frac{\partial \phi}{\partial y} + \varepsilon \frac{\partial \varepsilon}{\partial \phi_y} |\nabla \phi|^2 \right) + \frac{\partial}{\partial z} \left(\varepsilon^2 \frac{\partial \phi}{\partial z} + \varepsilon \frac{\partial \varepsilon}{\partial \phi_z} |\nabla \phi|^2 \right) + 4W\phi \left(-\phi \right) \left\{ \phi - \frac{1}{2} + \beta + a\chi \right\} \right]$$

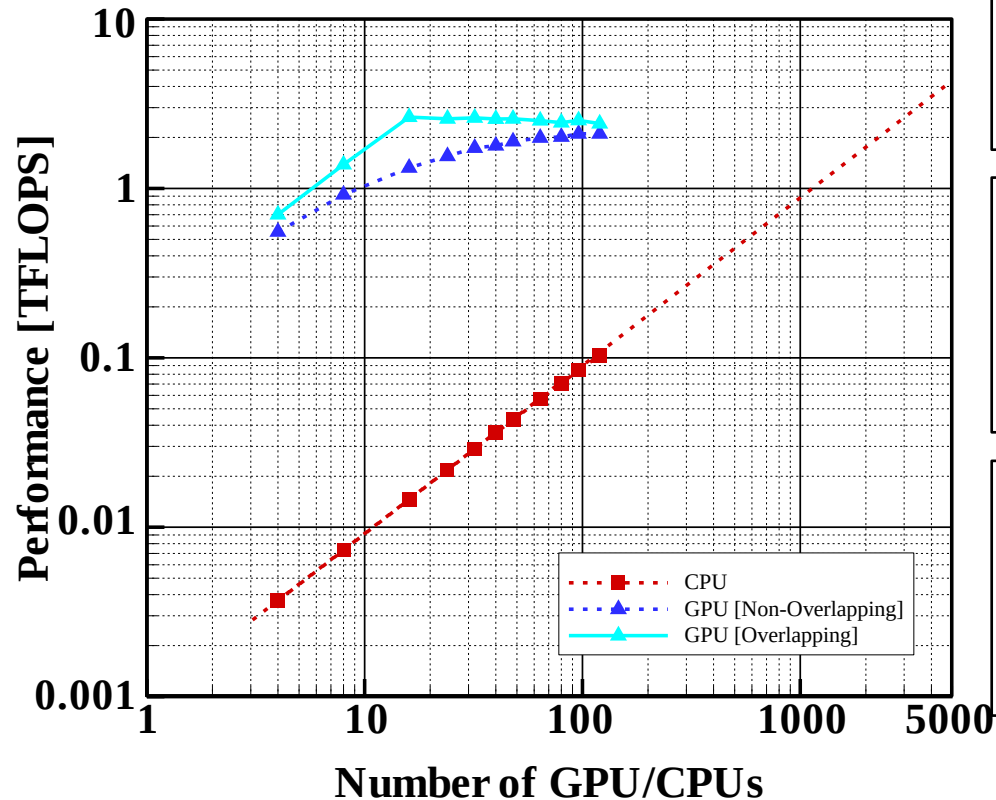
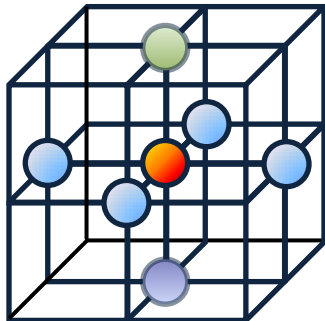
Thermal Conduction

$$\frac{\partial T}{\partial t} = k \nabla^2 T + \frac{L}{C} 30\phi^2 \left(-\phi \right) \frac{\partial \phi}{\partial t}$$

ϕ : 19 points to solve



T : 7 points to solve



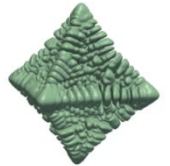
$t = 3.44 \times 10^{-8}$



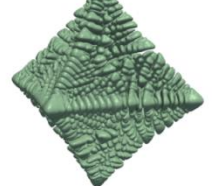
$t = 6.88 \times 10^{-8}$



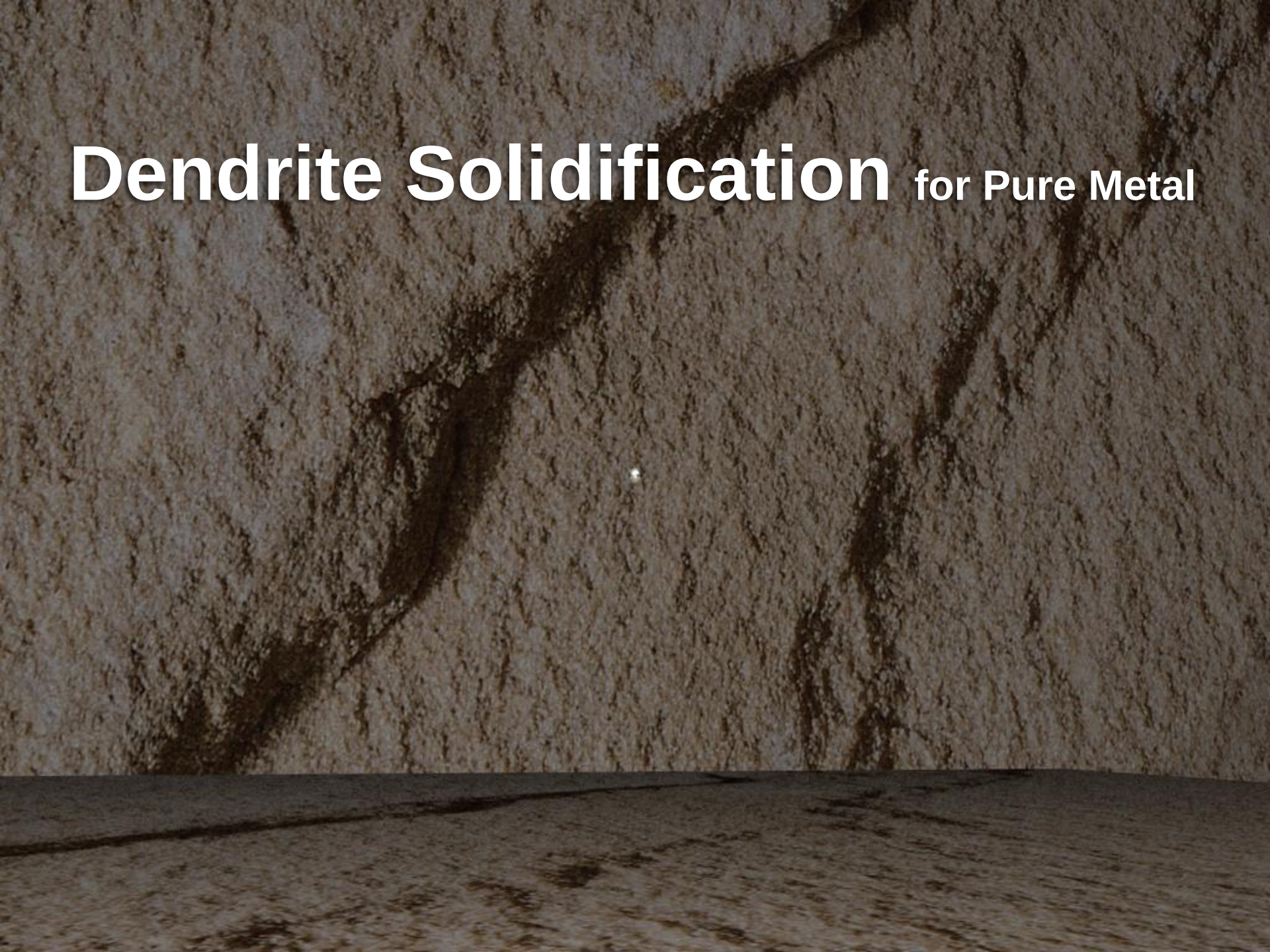
$t = 1.38 \times 10^{-7}$



$t = 2.06 \times 10^{-7}$



Dendrite Solidification for Pure Metal



Pulmonary Airflow Study

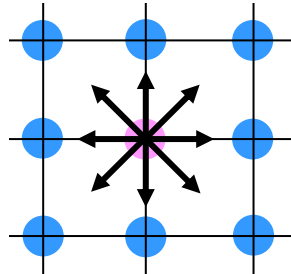
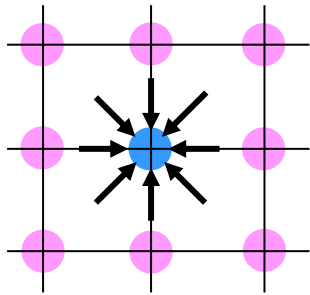
Lattice Boltzmann Method

$$\frac{\partial f_i}{\partial t} + \mathbf{e}_i \cdot \nabla f_i = -\frac{1}{\lambda} (f_i - f_i^{eq}) \quad f_i^{eq} = \rho w_i \left[1 + \frac{3}{c^2} \mathbf{e}_i \cdot \mathbf{u} + \frac{9}{2c^4} (\mathbf{e}_i \cdot \mathbf{u})^2 - \frac{3}{2c^2} (\mathbf{u} \cdot \mathbf{u}) \right]$$

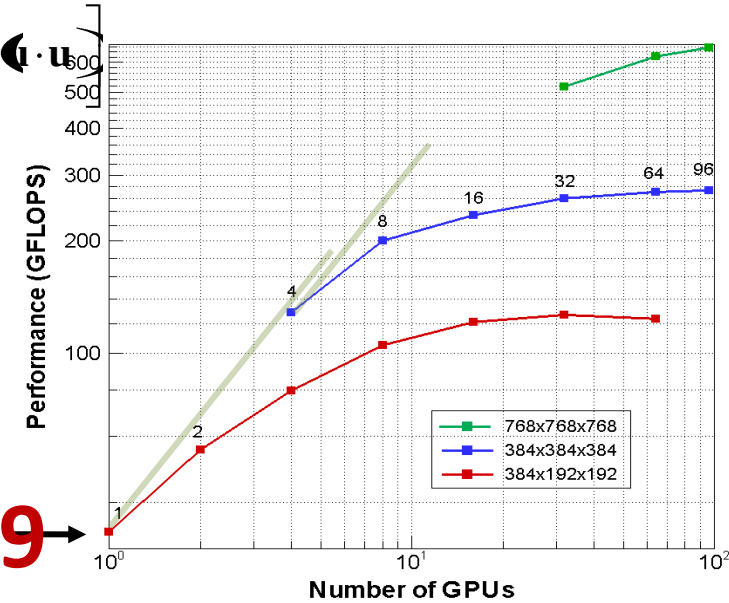
Collision step:

$$f_i(\mathbf{x} + \mathbf{e}_i \Delta t, t + \Delta t) = \bar{f}_i(\mathbf{x}, t) \quad \bar{f}_i(\mathbf{x}, t) = f_i(\mathbf{x}, t) - \frac{1}{\tau} (f_i(\mathbf{x}, t) - f_i^{eq}(\mathbf{x}, t))$$

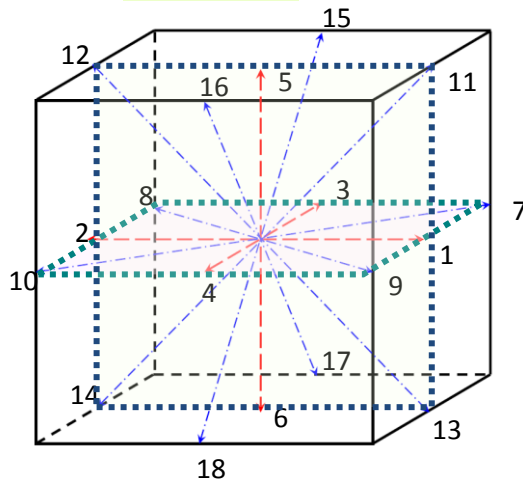
Streaming step:



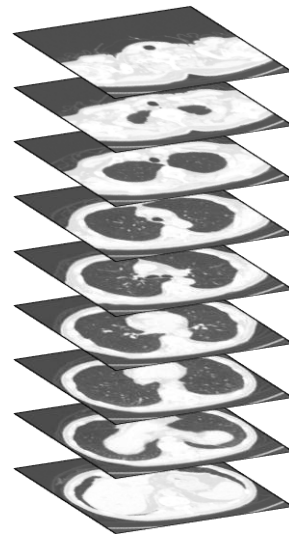
Prof. Takayuki Aoki
Collaboration with Tohoku University



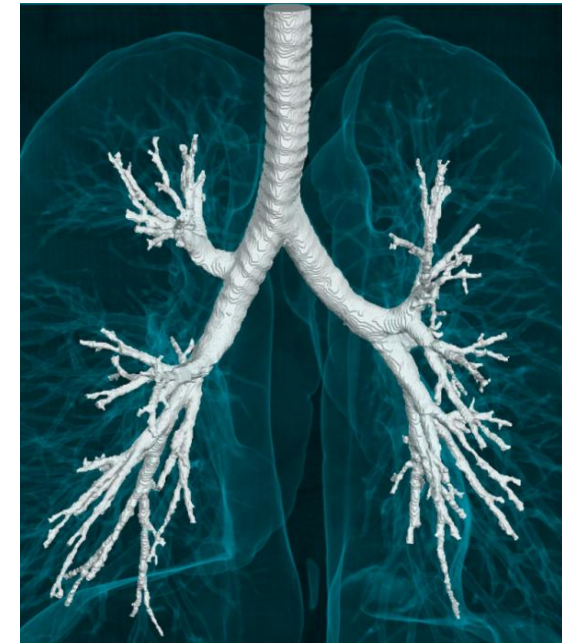
D3Q19



X-Ray CT images Airway structure Extraction

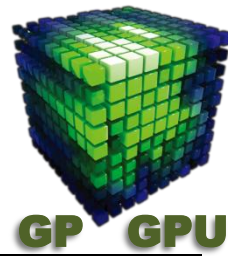


Lattice Boltzmann GPU computing

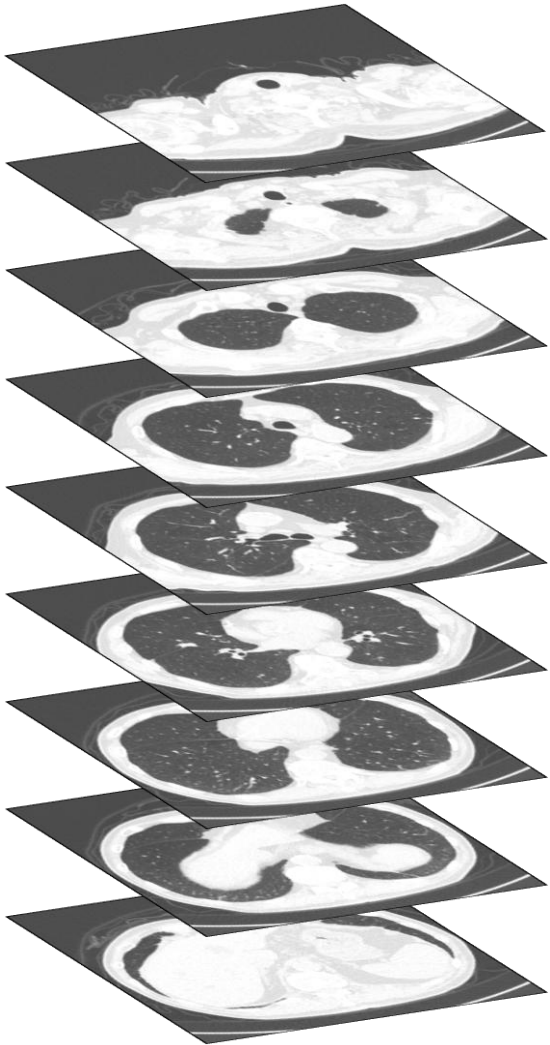


Pulmonary Airflow Study

Collaboration with Tohoku University



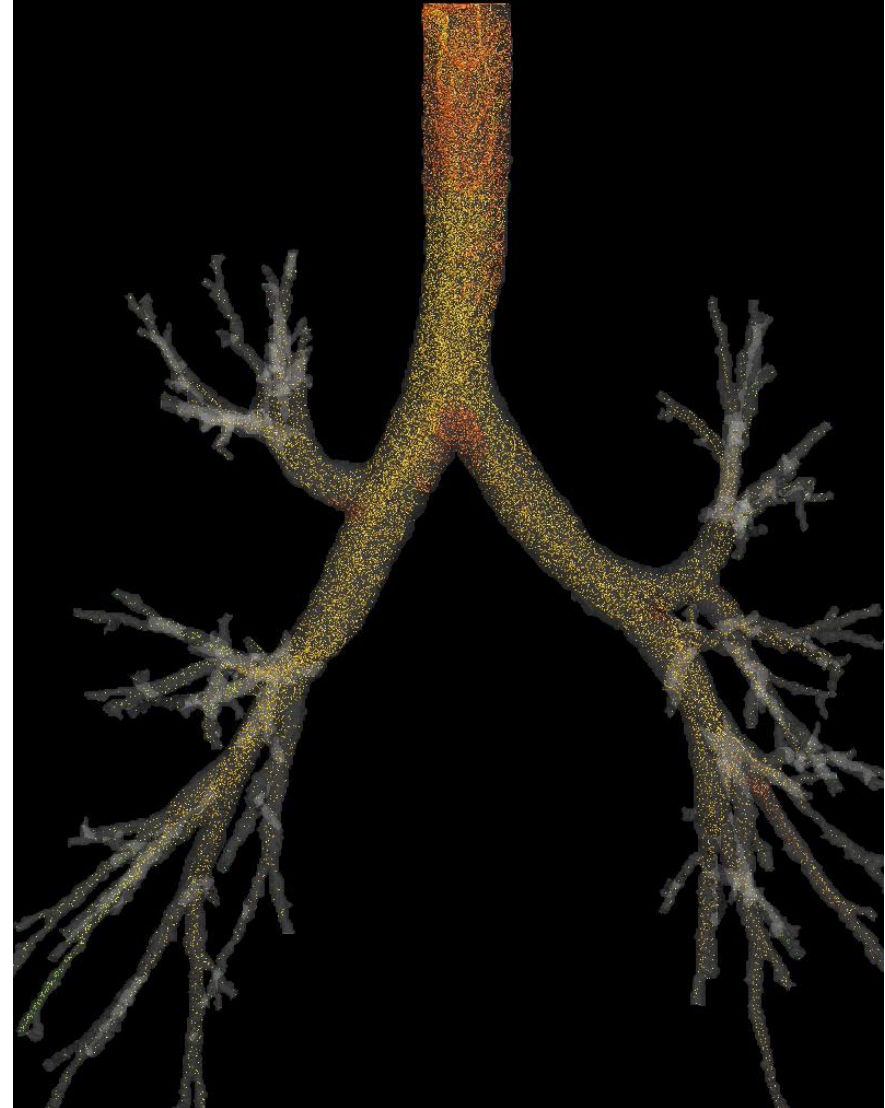
X-Ray CT images



Airway structure
Extraction



Lattice Boltzmann
GPU computing



Next Generation Numerical Weather Prediction[SC10]

Collaboration: Japan Meteorological Agency

Meso-scale Atmosphere Model:
Cloud Resolving Non-hydrostatic model
[Shimokawabe et. al. SC10 BSP Finalist]
ex. WRF(Weather Research and Forecast)



Typhoon ~ 1000km
1~ 10km
Tornado,
Down burst,
Heavy Rain

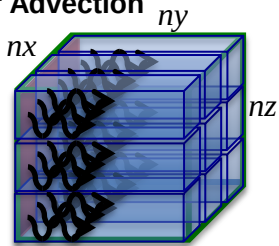
WSM5 (WRF Single Moment 5-tracer) Microphysics*
Represents condensation, precipitation and thermodynamic effects of latent heat release
1 % of lines of code, 25 % of elapsed time
⇒ 20 x boost in microphysics (1.2 - 1.3 x overall improvement)

ASUCA : full GPU Implementation

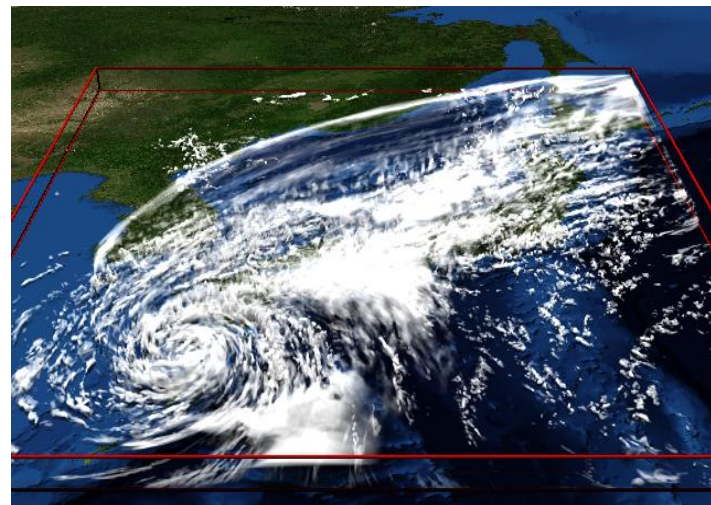
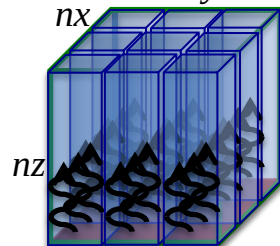
developed by Japan Meteorological Agency

TSUBAME 2.0 : 145 Tflops
World Record !!!

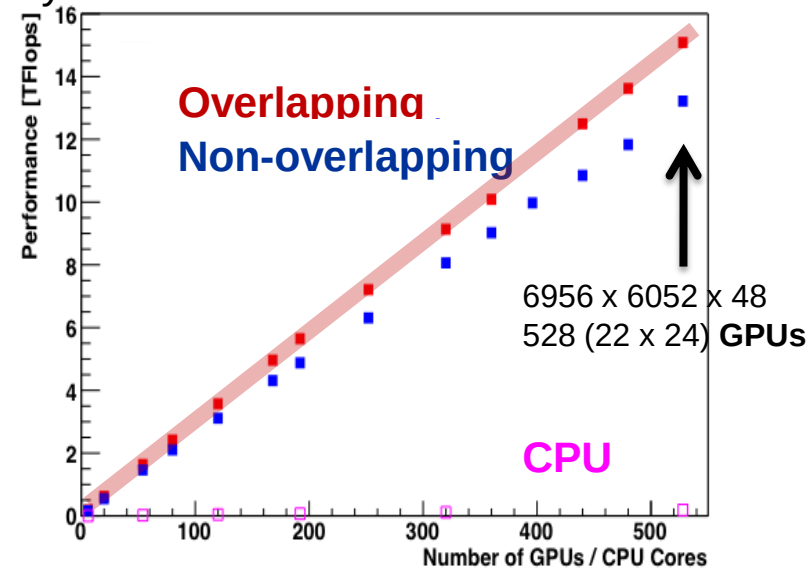
Block Division
for Advection



for 1-D
Helmholtz eq. ny



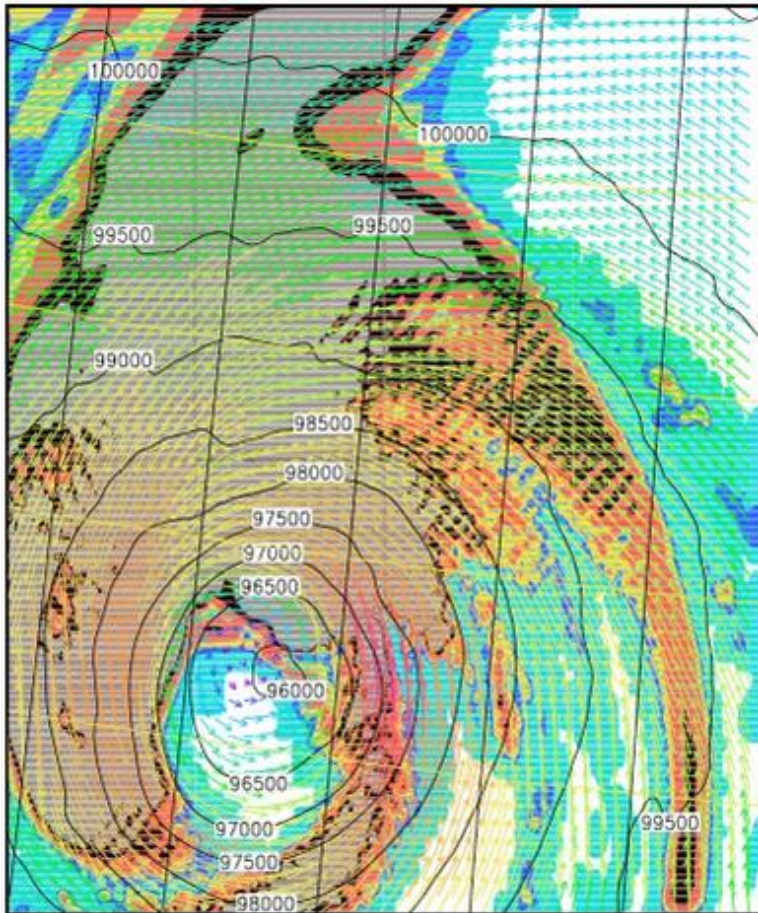
Typhoon



ASUCA Typhoon Simulation

2km mesh 3164×3028×48

uv and smqr T=1



70 minutes wallclock time
for 6 hour simulation time
(x5 faster)

Expect real-time with
0.5km mesh (for Clouds)
on TSUBAME2.0

MUPHY : Multi-Physics Simulator

[Massimo Bernaschi]

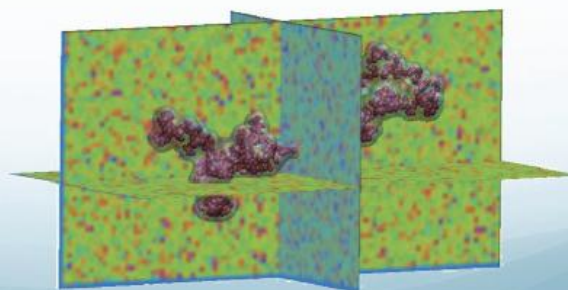
Broad spectrum of **fluid-particle** coupling mechanisms

Library of particle and molecular representations

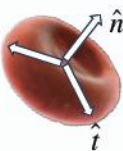
Library of fluid types

Co-polymers, colloidal suspensions, gels, biofluids, ...

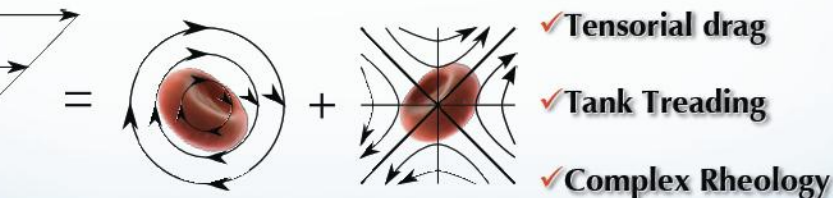
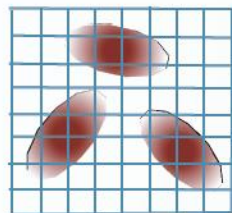
Multi-Platform



Ellipsoidal Suspensions of RBC



$$\delta_{\vec{v}_i \vec{v}_j}^{\lambda} (x) = \prod_{\alpha=x,y,z} \delta_{\vec{v}_i \alpha}^{\lambda} ((Qx)_{\alpha})$$



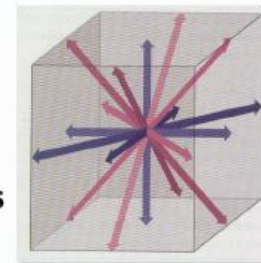
Strip off DOFs: O(10) per RBC

Blood Plasma via Lattice Boltzmann

From (minimal) Boltzmann equation

$$(\partial_t + \vec{v} \partial_x) f(\vec{x}, \vec{v}, t) = -\omega(f - f^{eq})(\vec{x}, \vec{v}, t)$$

Collision + Streaming of a set of discrete velocities



Superset of the Navier-Stokes dynamics

Exact streaming (no self-convected): uniform mesh

Complexity O(N)

Enable complex geometries

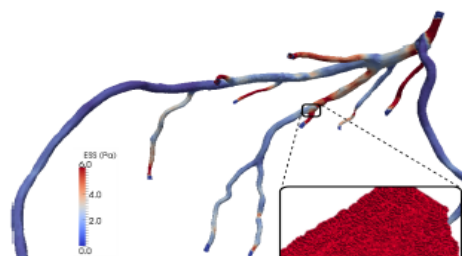
MUPHY Performance on the TiTech GPU cluster

The performance of 1536 GPUs for the Lattice Boltzmann kernel is the same as the whole Jülich Bluegene/P system (294712 cores!)

# of GPUs	time	efficiency
256	76.16	N.A.
512	38.52	98.86%
1024	19.95	95.37%
1536	13.43	94.43%

x3 40 rack BG/P

# of GPUs	time
256	648.23
512	327.97



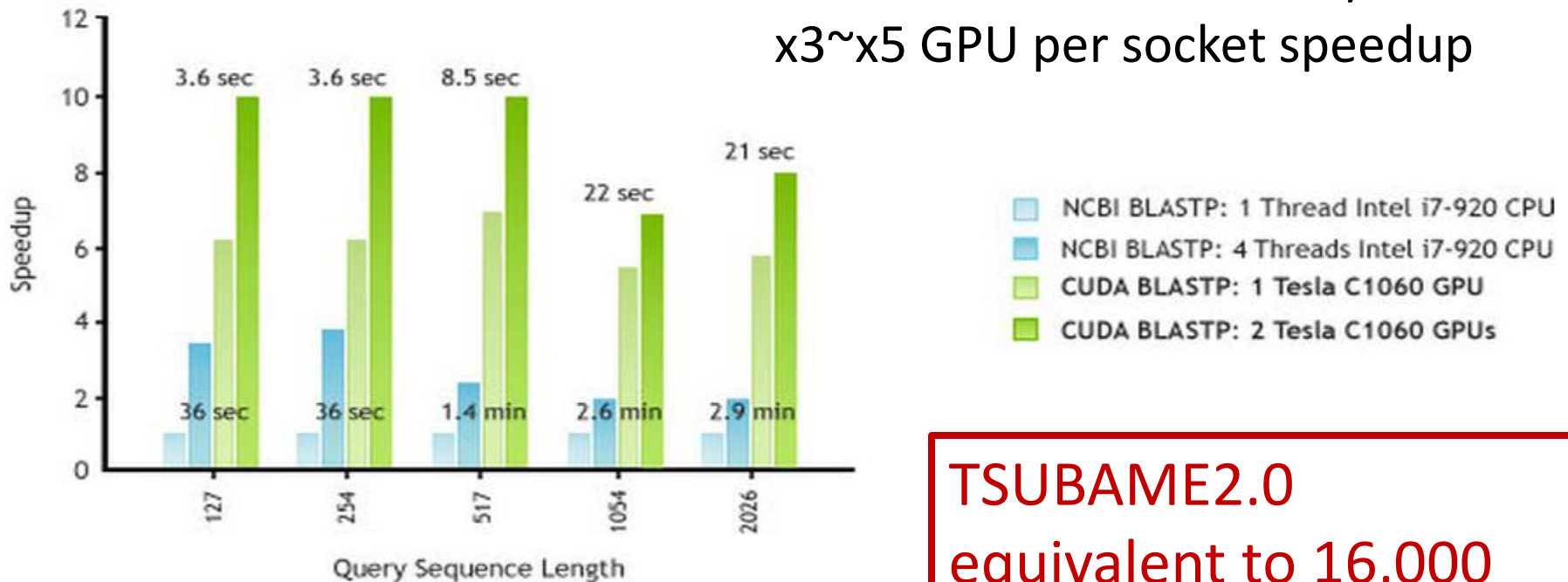
LB and Molecular Dynamics
 1×10^9 nodes, 100×10^6 cells

Bioinformatics: BLAST on GPUs

- CUDA-BLASTP (NTU)

http://www.nvidia.com/object/blastp_on_tesla.html

CUDA-BLASTP vs NCBI BLASTP Speedups



Data Courtesy of Nanyang Technological University, Singapore

Previous-Generation CPU/GPU
x3~x5 GPU per socket speedup

TSUBAME2.0
equivalent to 16,000
sockets (100,000 cores)

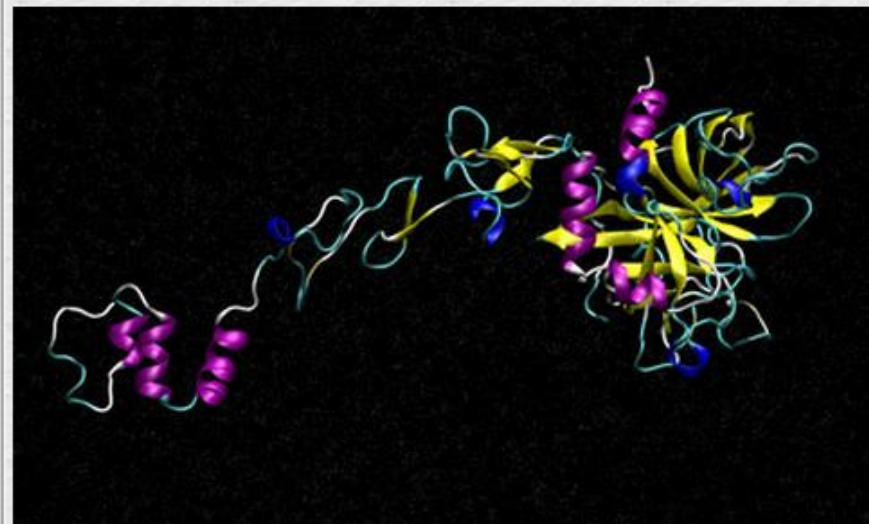
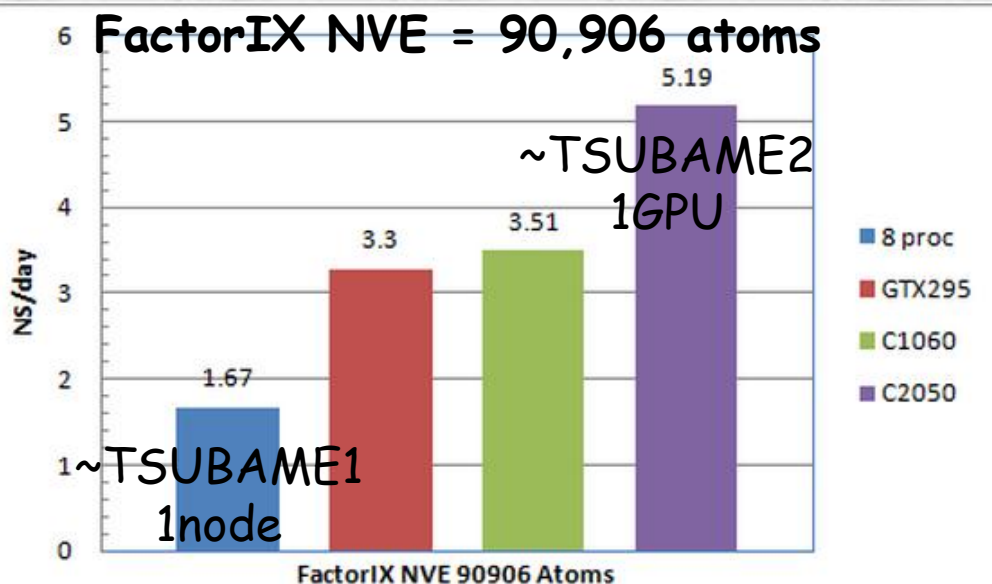
- GPU-BLAST (CMU)

<http://eudoxus.cheme.cmu.edu/gpublast/gpublast.html>

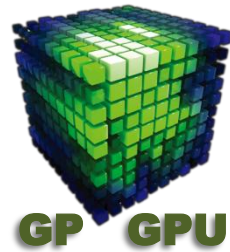
“4 times speedup on Fermi GPUs”

GPU Enabled Amber 11

- 1GPU $\approx$ 8 Core CPU node x 4~20
 - Already in service on TSUBAME2.0
 - Waiting for Multi-GPU version presented at GTC 2010...
- Equivalent to 300,000 CPU cores on TSUBAME2.0
 - Assuming x10 speedup per socket



Electric Power Consumption for Multiple-GPU Applications



GP GPU

CFD Applications

Comparing to TSUBAME Opteron CPU 1 core,



MAX 200W/GPU

Weather Prediction: x80 (10 GPU = 1000 CPU)

Lattice Boltzmann: ~ x100

Phase Field Model: x170

When we assume the acceleration is typically x100,
1 GPU < 200W for our applications

100 GPU : 20 kW

10000 CPU core of **TSUBAME 1.2** : 1MW (1000kW)



TSUBAME ~
1MW/10000CPU

1/50 Ultra low Power



But it is not just HW Design...SOFTWARE(!)

- The *entire software stack* must preserve bandwidth, hide latency, lower power, heighten reliability for Petascaling
- Example: TSUBAME1.2, Inter-node GPU $\Leftrightarrow$ GPU achieves only 100-200MB/s in real apps
 - c.f. 1-2GB/s HW dual rail IB capability
 - Overhead due to unpinnable buffer memory?
 - New Mellanox driver will partially resolve this?
 - Still need programming model for GPU $\Leftrightarrow$ GPU
- Our SW research as CUDA CoE (and other projects such as Ultra Low Power HPC)

Tokyo Tech CCOE

1. Auto-Tuning GPU Numerical Libraries

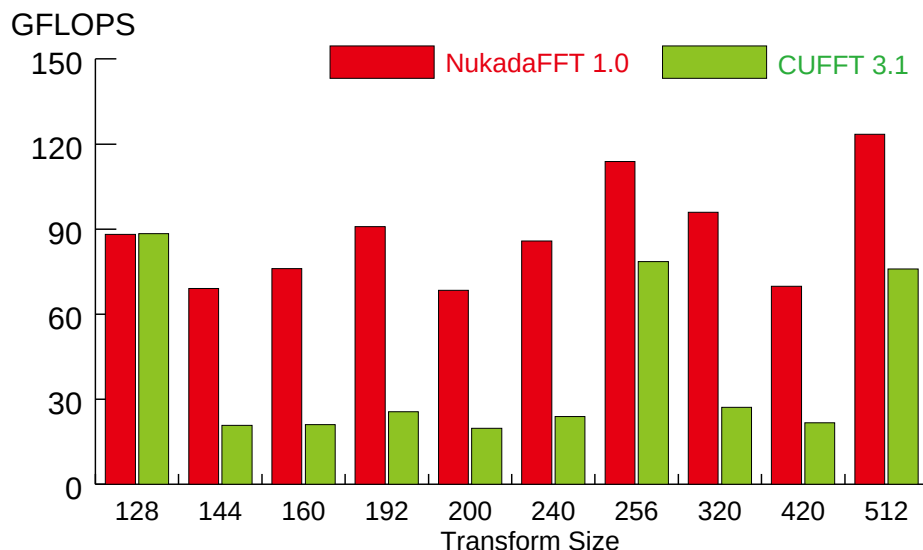


NukadaFFT 1.0 release !? [sc08,sc09]

NukadaFFT library is a very fast auto-tuning FFT library for CUDA GPUs.

Tuning Parameters:

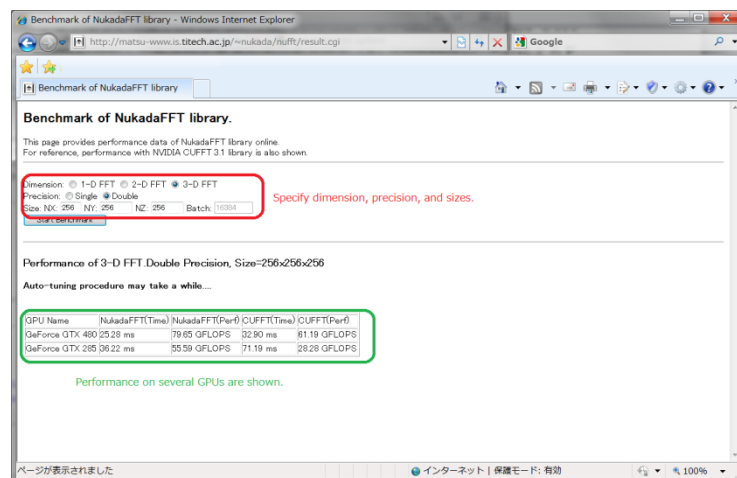
- (1) Factorization of transform size.
 - (2) # of threads launched per SM.
 - (3) Padding insertion patterns for shared memory
- The library generates many kinds of FFT kernels and selects the fastest one. (exhaustive)



Performance of 1-D FFT.

(Double Precision, batch=32,768, GPU=GeForce GTX 480.)

For more details, you can see GTC 2010 Research Poster, or catch the author.



You can try online benchmarking as for the size you are interested in.

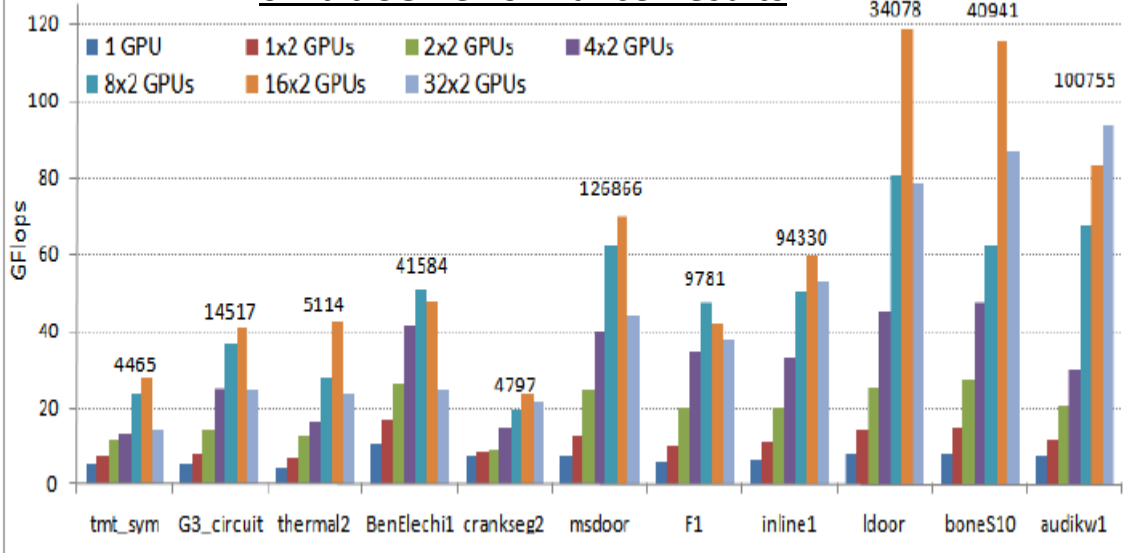
<http://matsu-www.is.titech.ac.jp/~nukada/nufft/>

“Auto-Tuning” Sparse Iterative Solvers on a Multi-GPU Clusters

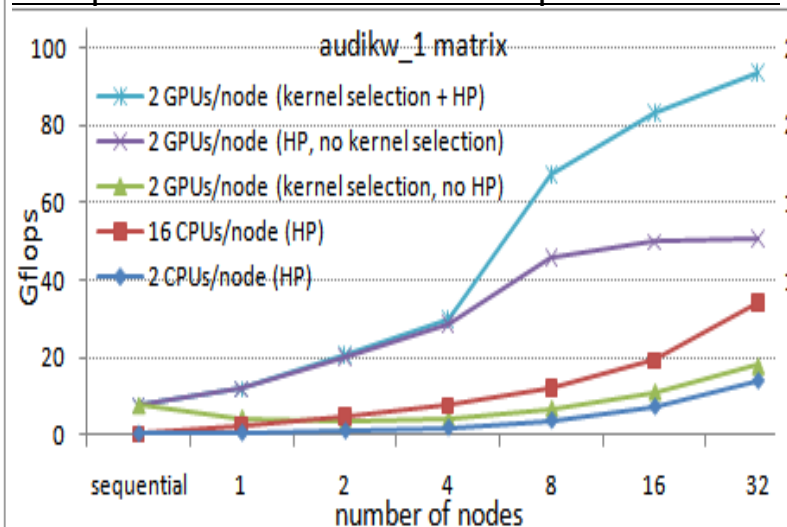
(A. Cevahir, A. Nukada, S. Matsuoka)

- High communication due to irregular sparsity
 - GPU computing is fast, communication bottleneck is more severe
- We propose techniques to achieve scalable parallelization
 - A new sparse matrix-vector multiplication (MxV) kernel [*ICCS'09*]
 - **Auto-selection of optimal running MxV kernel** [*ISC'10*]
 - Minimization of communication between GPUs [*ISC'10*]
- Two methods: **CG and PageRank (on 100 CPUs)**

64-bit CG Performance Results



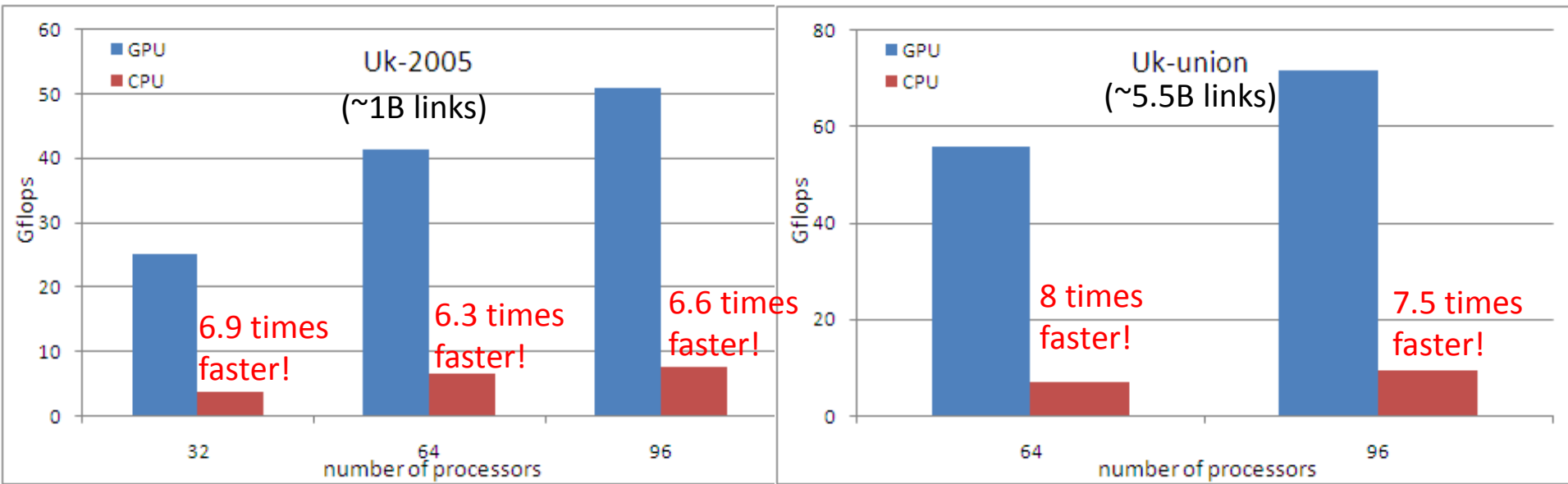
Comparisons of different CG implementations



PageRank Results

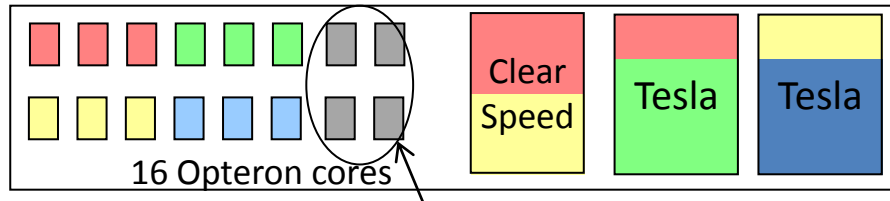
(A. Cevahir, A. Nukada, S. Matsuoka)

- Huge data is processed -- bigger data scales better
- A state-of-the-art PageRank method is implemented
(based on *Cevahir, Aykanat, Turk, Cambazoglu: IEEE TPDS June 2010*)
(Problem bandwidth restricted – socket - to - socket comparison)



High performance Linpack on heterogeneous Supercomputers [IPDPS 08,10]

- Dense linear computation on large heterogeneous systems
 - All types of processors are used for kernel (DGEMM)
 - Even supporting “heterogeneous nodes”
 - Load balancing techniques, tuning for reducing PCIe communication

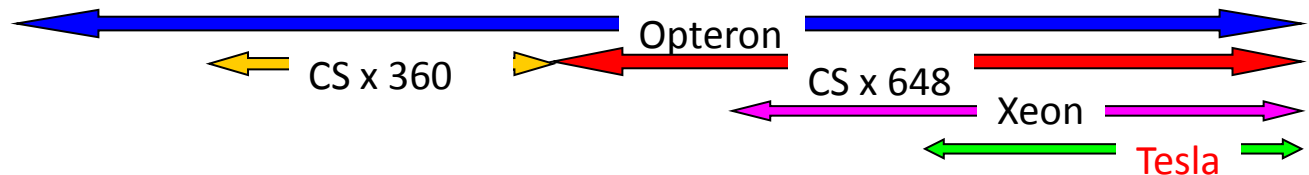


Mapping between MPI processes and heterogeneous processors

Some cores are dedicated for PCI comm

- >10,000 CPU cores, >600 Teslas, >600 ClearSpeed are cooperatively used
 - 87.01TFlops, 53% efficiency → 7 times improvements on Top 500 ranking

	Jun06	Nov06	Jun07	Nov07	Jun08	Nov08	Jun09
Linpack speed	38.18TF	47.38	48.88	56.43	67.70	77.48	87.01
Rank	7	9	14	16	24	29	41



Tokyo Tech CCOE

2. GPU Power Performance
Modeling and Heterogeneous
Scheduling

Measuring GPU Power Consumption

- Two power sources
 - Via PCI Express: < 75 W
 - Direct inputs from PSU: < 240 W
- Uses current clamp sensors around the power inputs

Precise and Accurate Measurement with Current Probes

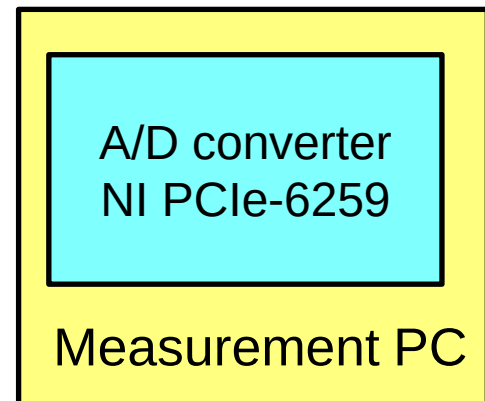


Attaches current sensors to two power lines in PCIe

+



Direct power inputs from PSU



Reads currents at 100 us interval

Statistical Modeling of GPU Power Consumption [IEEE HPPAC10]

- Regularized linear regression
 - Finds linear correlation between per-second PMC values and average power of a kernel execution
 - Aggregates 3 kinds of global memory loads/stores
 - gld: $\text{gld_32b} + \text{gld_64b} * 2 + \text{gld_128b} * 4$
 - Regularization for avoiding overfitting to training data (Ridge Regression [10])

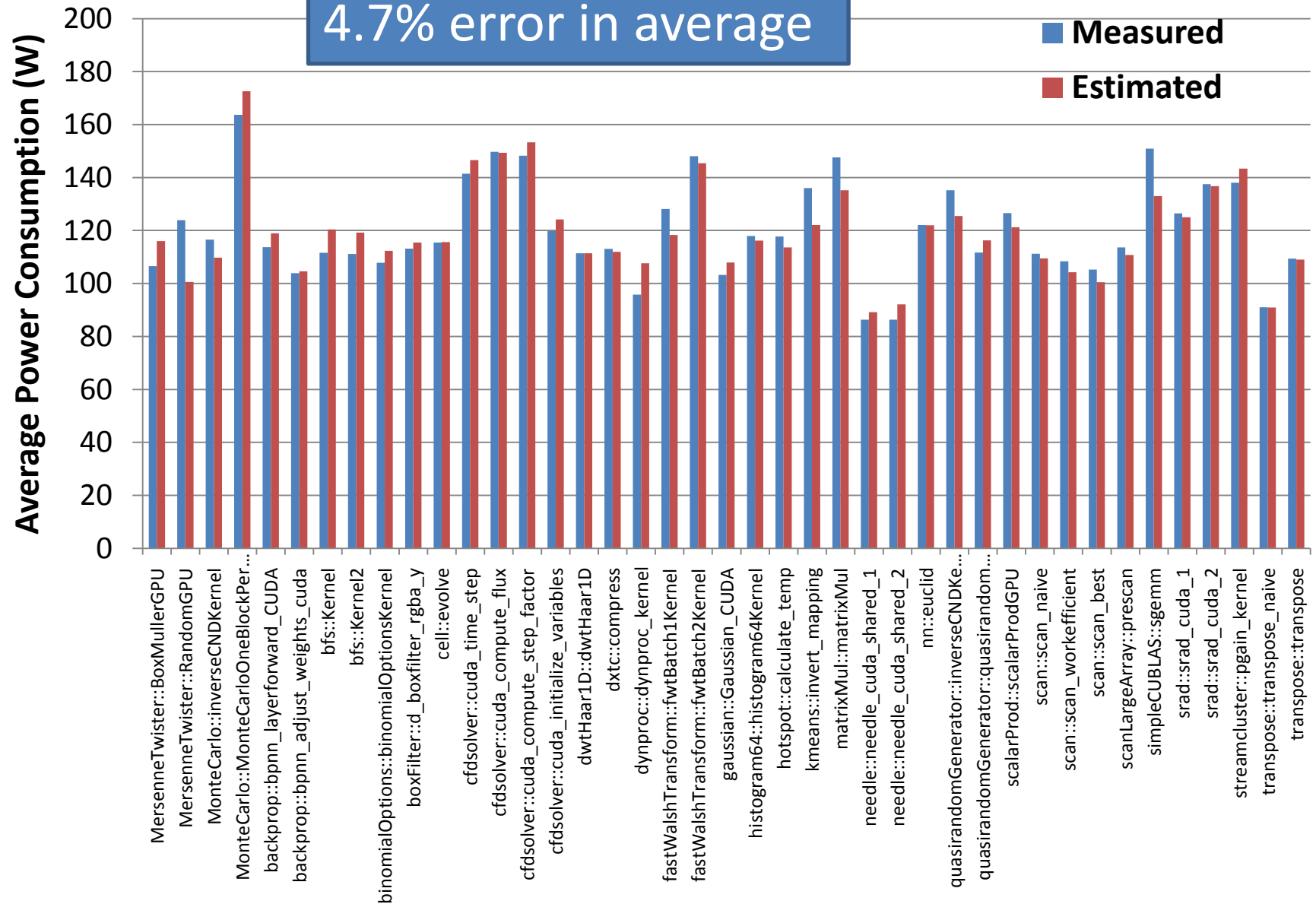
Average power of a kernel (Watts)

$$p = \sum_{i=1}^n \alpha_i c_i + \beta$$

Per-second PMC values

Actual vs Estimated Power

4.7% error in average



Low Power Scheduling in GPU Clusters

[IEEE IPDPS-HPPAC 09]

Objective : Optimize CPU/GPU

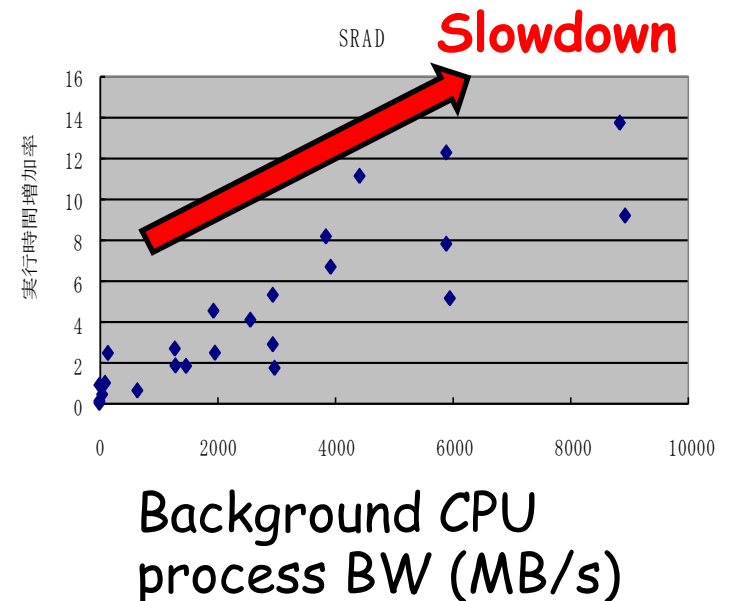
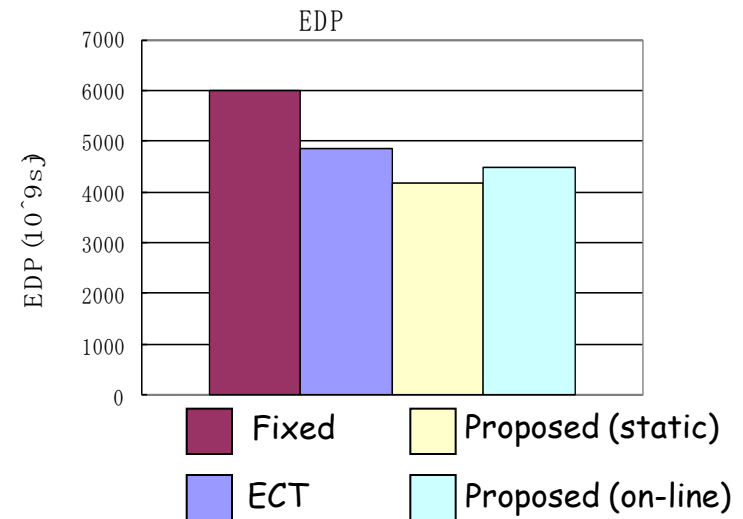
Heterogeneous

- Optimally schedule mixed sets of jobs executable on either CPU or GPU but w/different performance & power
- Assume GPU accel. factor (%) known

30% Improvement Energy-Delay Product

TODO: More realistic environment

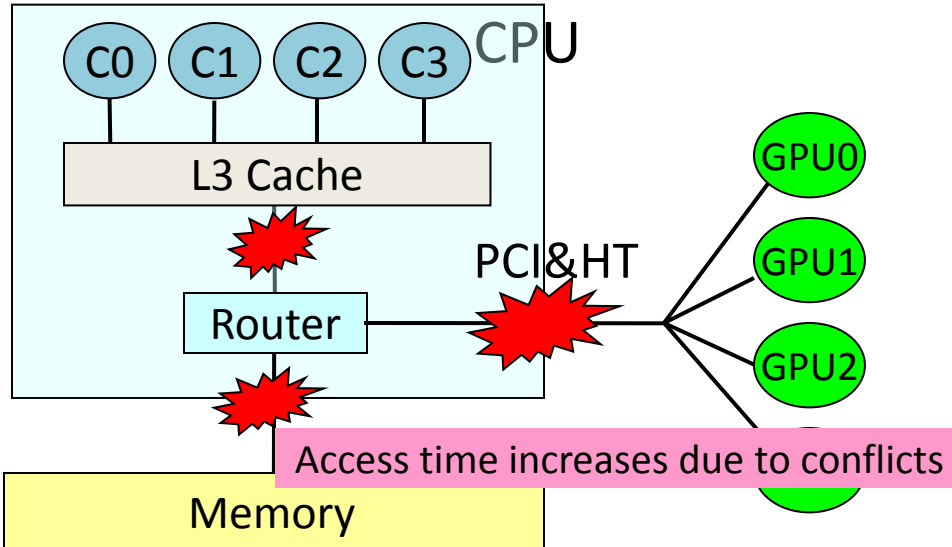
- Different app. power profile
- PCI bus vs. memory conflict
 - GPU applications slow down by 10 % or more when co-scheduled with memory-intensive CPU app.



Intelligent Task Scheduling for GPU clusters

- considering conflicts on memory bus and PCI-E bus -

4 processes share a node.



Assuming the following information of each job is known

- Execution time (at stand alone case)
- Total memory access (from performance counter)
- Total PCI-E transfer (from CUDA profiler)

Actual execution time is estimated by our model.

T. Hamano, et al. IPSJ SIGHPC-124 (domestic)

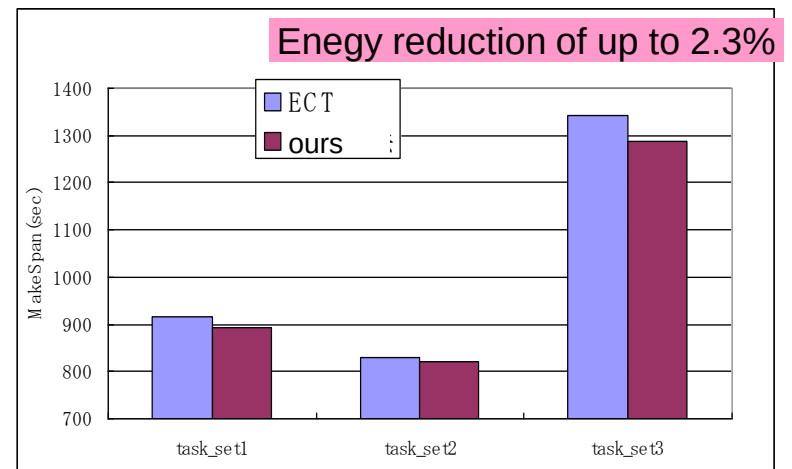
Status of each process is one of the following:

- (1) Memory access
- (2) PCI transfer (assume blocking cuMemcpy)
- (3) No memory access

Speed-down ratio of a process = r (percentage of each status) $\times \alpha$ (conflict coefficient)

$$\sum_{s_0} \sum_{s_1} \sum_{s_2} \sum_{s_3} \alpha(s_0, \{s_1, s_2, s_3\}) r_0(s_0) r_1(s_1) r_2(s_2) r_3(s_3)$$

We integrated this model into ECT(Earliest Complete Time) scheduling.



Tokyo Tech CCOE

3. GPU Fault Tolerance

NVCR : A Transparent Checkpoint/Restart for CUDA

CUDA Checkpoint using BLCR (Takizawa, 2009)

- Save all data on CUDA resource.
- Destroy all CUDA context.
- BLCR
- Reallocate CUDA resource.
- Restore data.

After reallocation, the address or handle may differ from previous one.

Handles are Opaque pointers : can be virtualized.

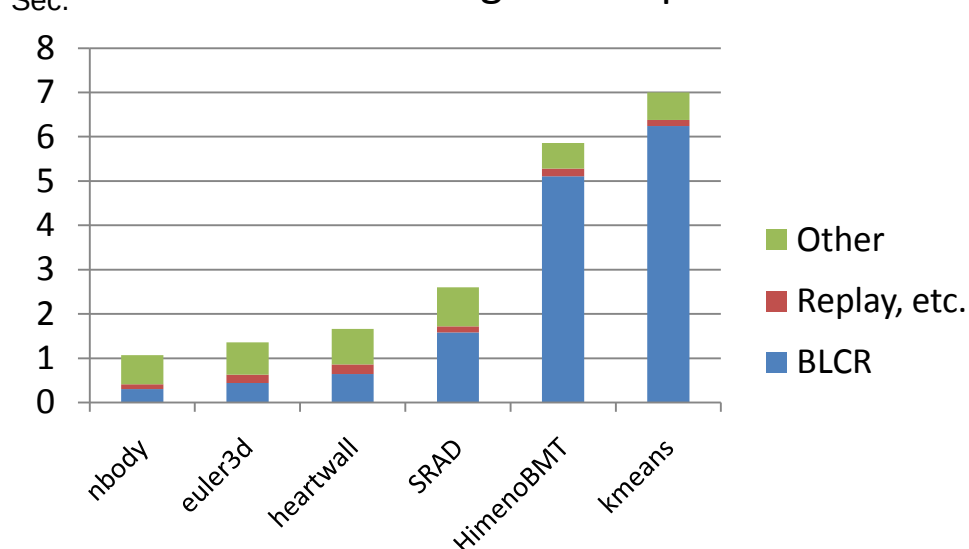
Addresses must be visible for user application.

Our answer is 'REPLAY'.

NVCR records calls of all APIs related to device memory address such as;
`cuMem*()`, `cuArray*()`, `cuModule*()`.

We added some optimization of reducing the API records. For example, locally resolved alloc-free pairs are removed from the record.

Sec. Overhead of single checkpoint



The time for Replaying is quite negligible compared with the main BLCR procedure to write image to HDD.

NVCR supports

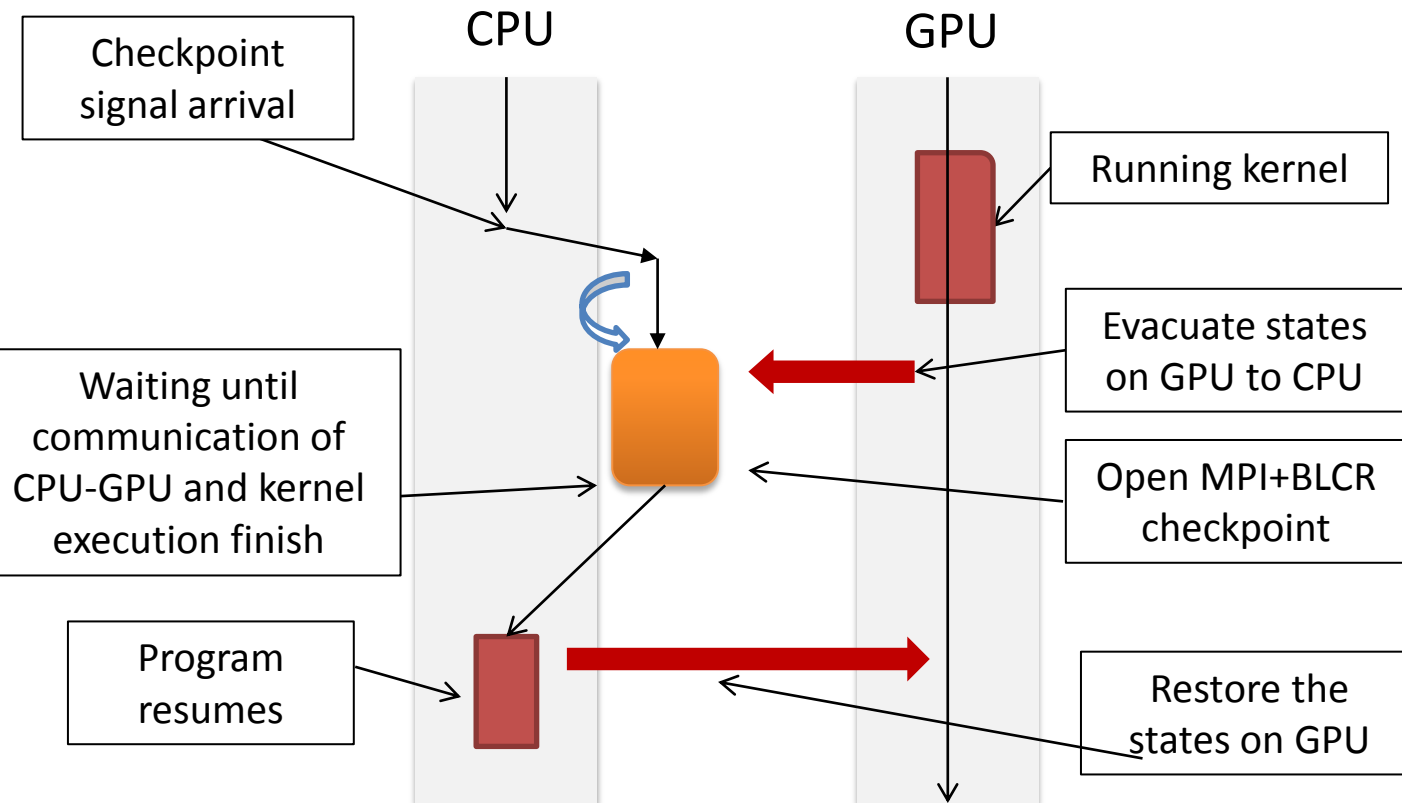
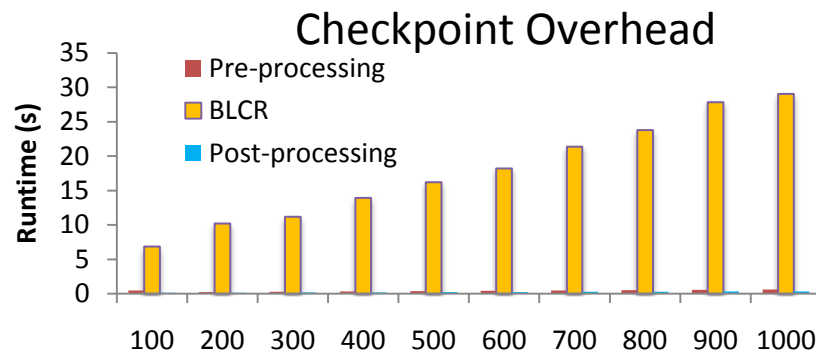
- All CUDA 3.0 APIs
- OpenMPI
- Both CUDA runtime and CUDA driver API.

CUDA pinned memory is supported partially ... Full support requires NVIDIA's official support.

MPI CUDA Checkpointing

Checkpoint for MPI+CUDA applications

- Very important in large-scale GPU systems
- Supports a majority of the CUDA RT API
- Based on BLCR and OpenMPI



Data size (MB)

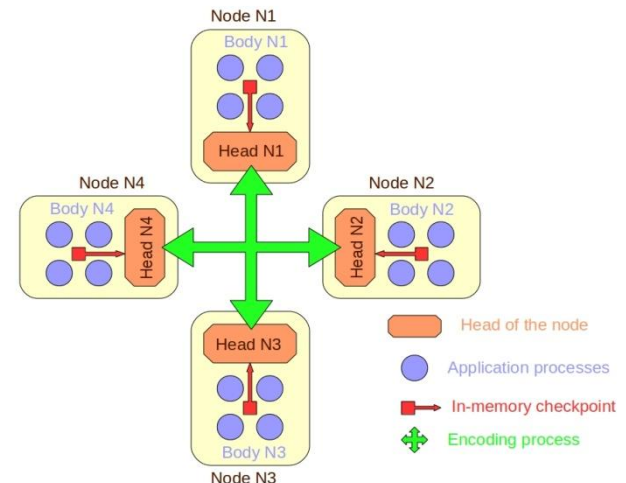
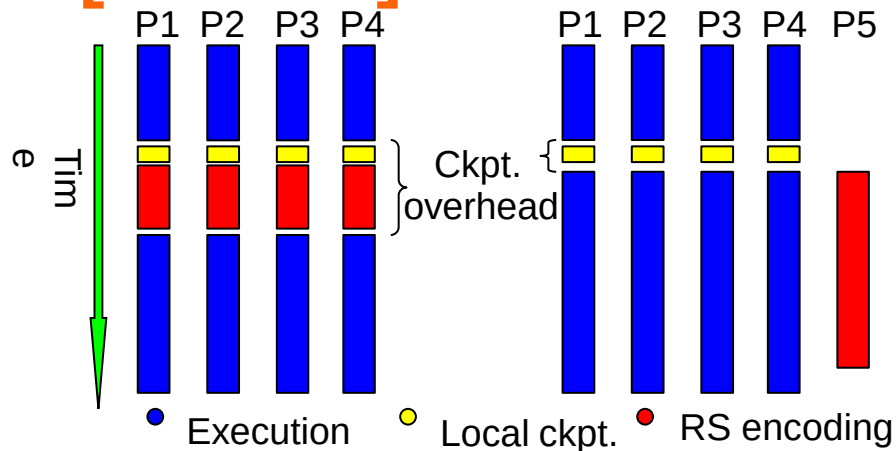
- GPU state saving and restoring has little overall performance impact.
- Can be significantly faster with our scalable checkpointing algorithms

**More details at
the Research
Poster Session**

Hybrid Diskless Checkpoint

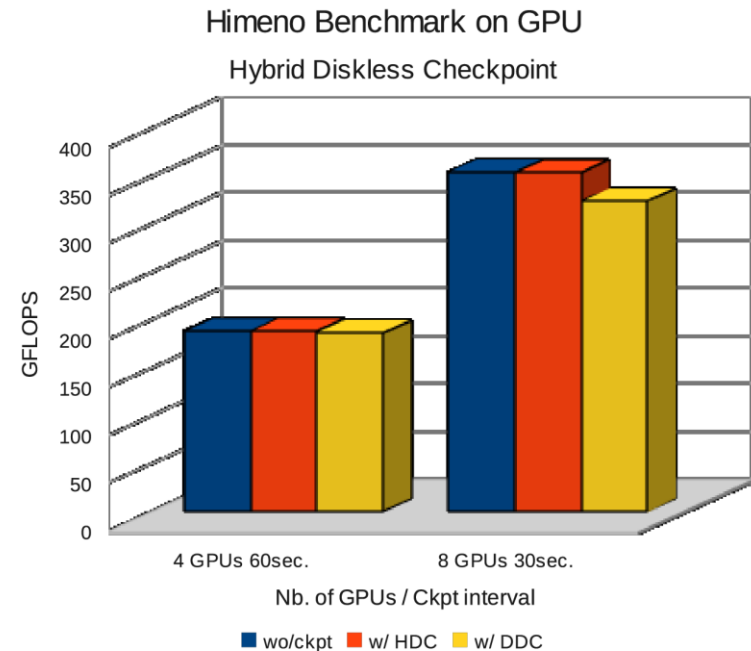
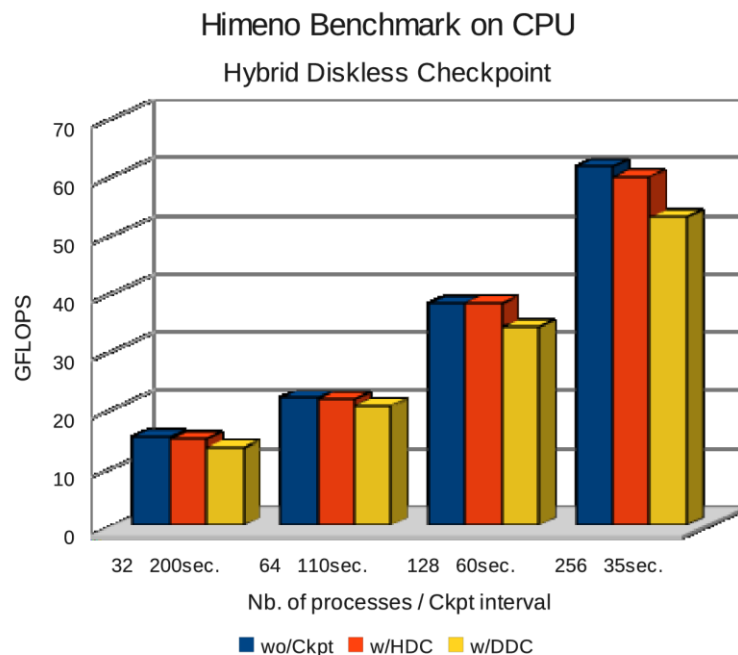
- **Problem:** Decrease the ckpt. overhead with erasure codes on large-scale GPU systems.
- Scalable encoding algorithm and efficient group mapping [CCGrid10]
- Fast encoding work on idle resource (CPU/GPU)

[HiPC10]



Hybrid Diskless Checkpoint

- Less than 3% of ckpt. overhead using idle resources



- We are currently combining this technique with our MPI CUDA Checkpoint.

Software Framework for GPGPU Memory FT

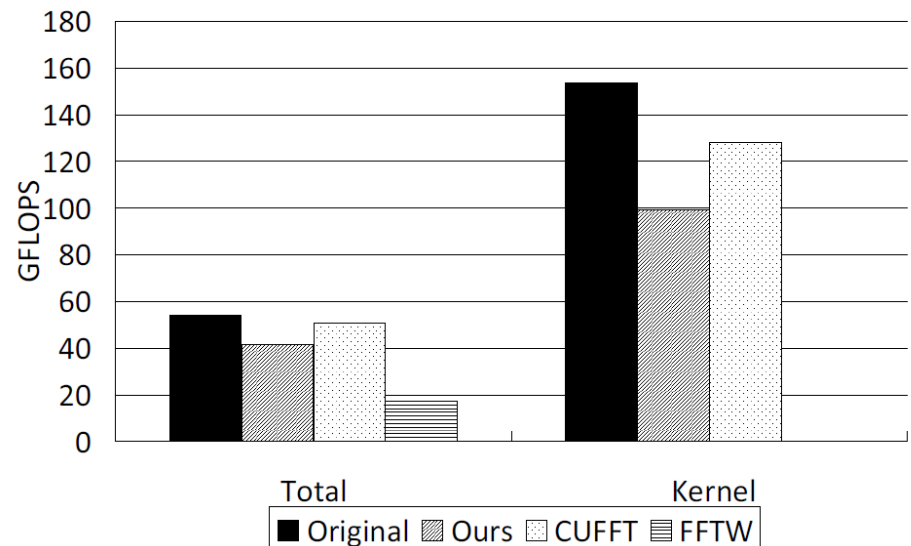
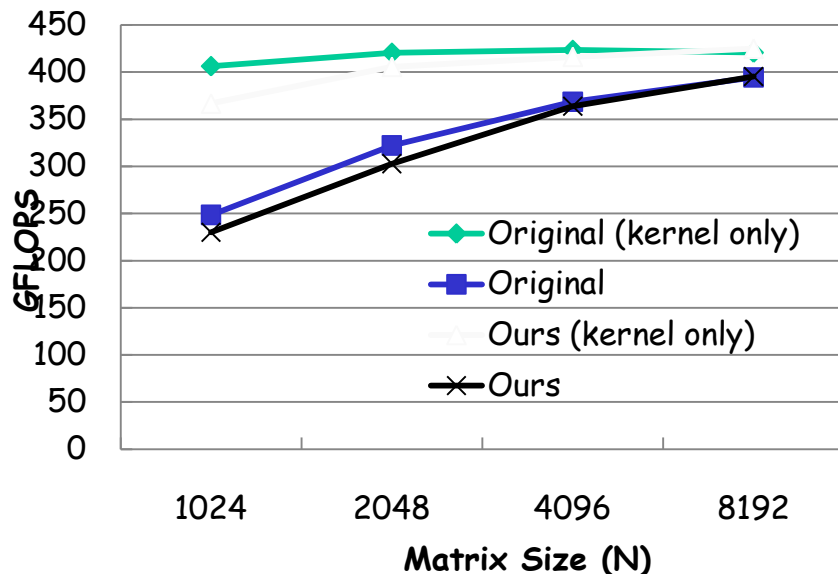
[IEEE IPDPS 2010]

- Error detection in CUDA global memory + Checkpoint/Restart
- Works with existing NVIDIA CUDA GPUs

Lightweight Error Detection

- *Cross Parity* for 128B blocks of data
- Detects a single-bit error in a 4B word
- Detects a two-bit error in a 128B block
- No on-the-fly correction → Rollback upon errors

Exploit the latency hiding capability of GPUs for data correctness

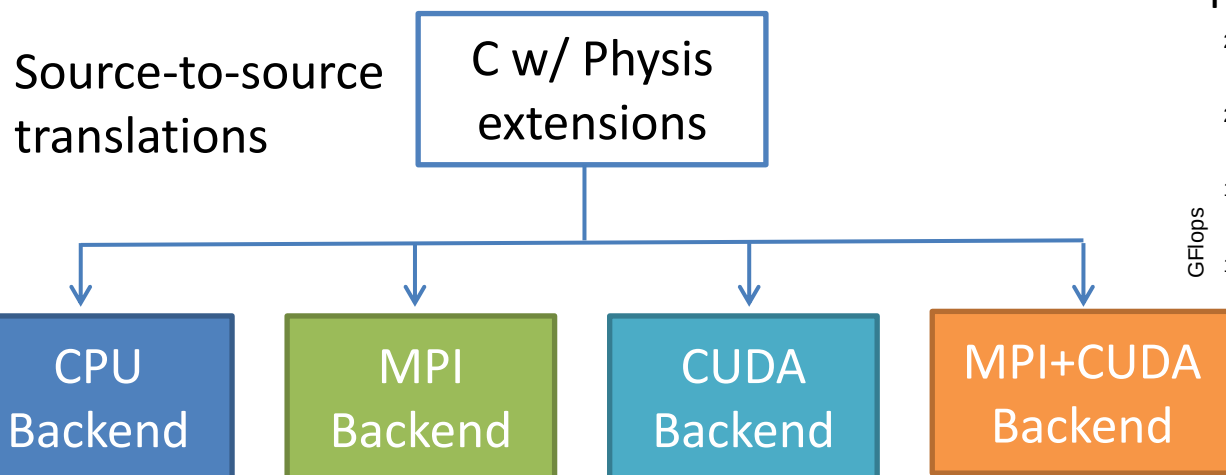


Tokyo Tech CCOE

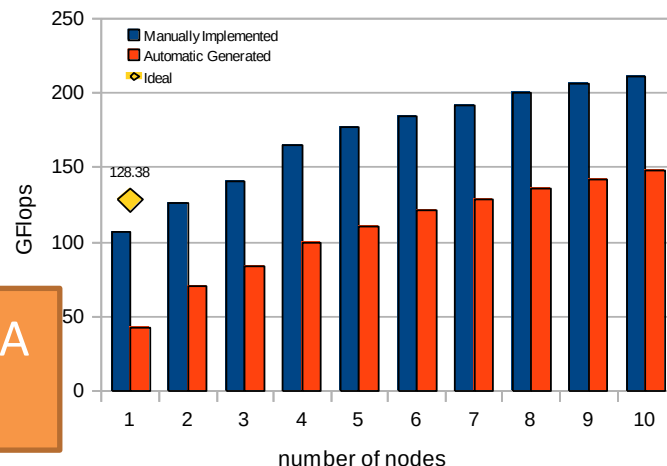
4. GPU Higher-Level
Programming Model

Physis: A Domain Specific Framework for Large-Scale GPU Clusters

- Portable programming environment for stencil computing on structured grids
 - Implicitly parallel
 - Distributed shared memory
- Provides concise, declarative abstractions for stencil computing



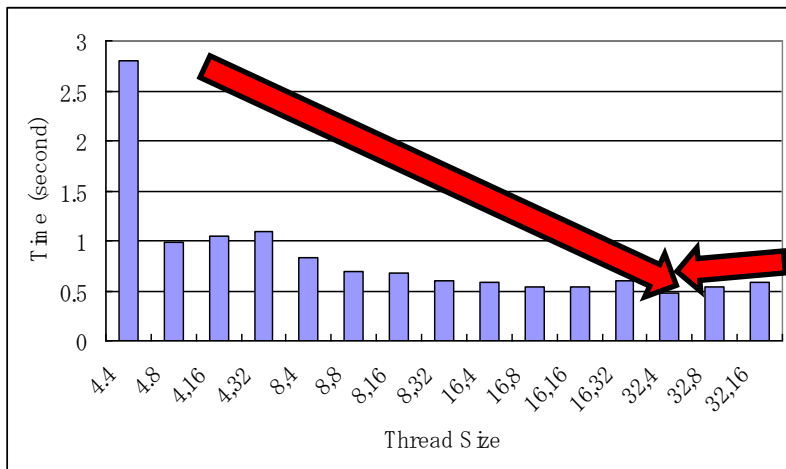
Preliminary Performance Results



Towards Auto-Tuning of Stencil Computation on GPU-accelerated Clusters

- Programming on GPU clusters gets harder because of
 - Hybrid programming with CUDA/MPI/threads
 - # of tuning parameters is increasing.
 - Hiding latency

→ Auto-tuning is more important



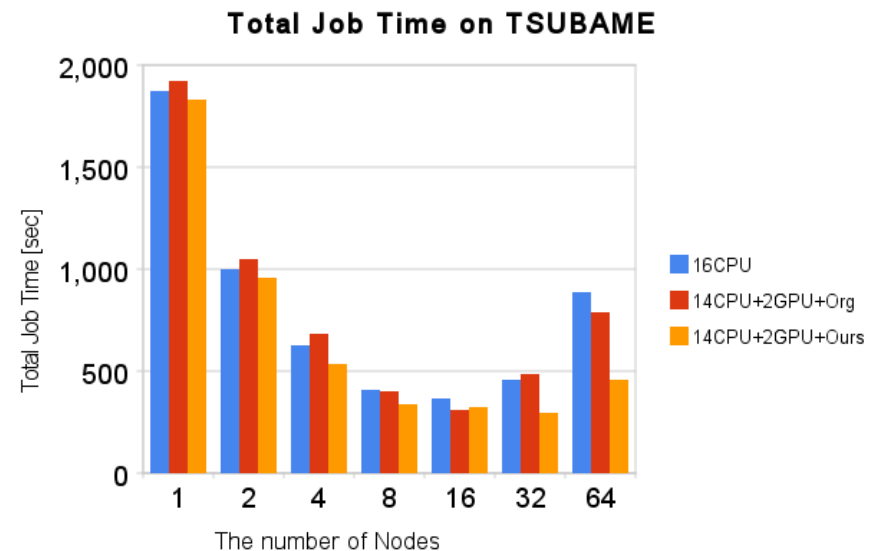
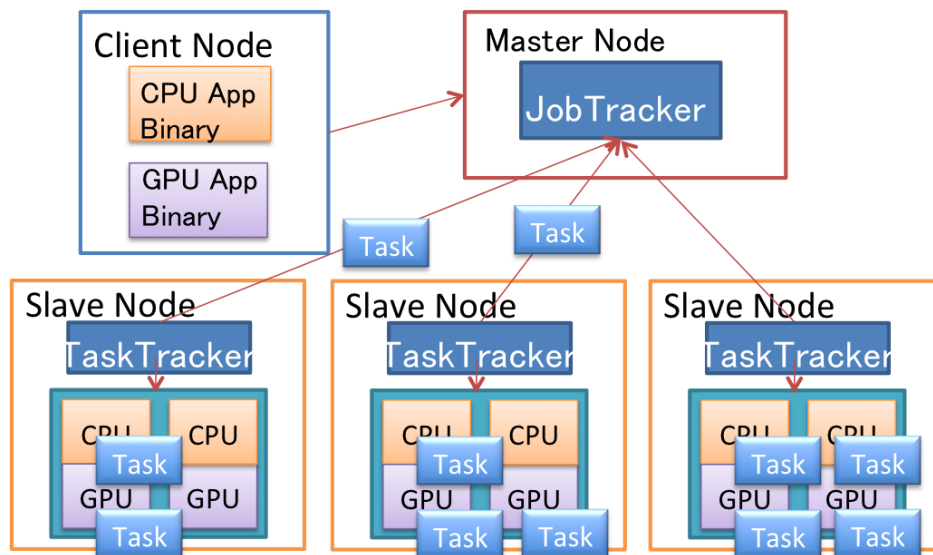
MapReduce

for GPU-based Heterogeneous Clusters

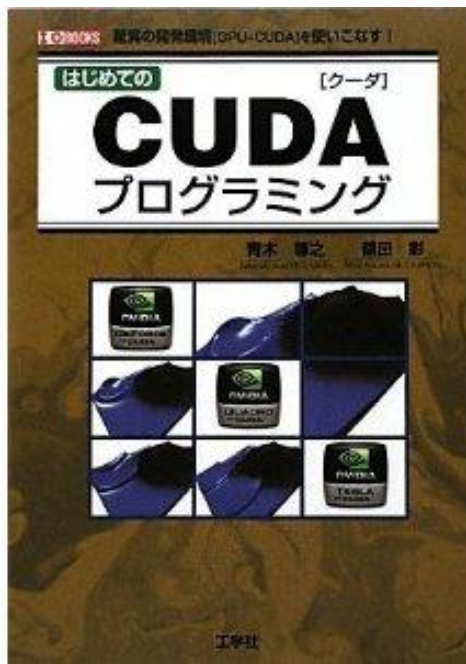
[MAPRED'2010]

Problem :

- Map Task Scheduling for efficient execution
 - Depends on **running task characteristics** and **underlying environments**
- Hybrid Map Task Scheduling
 - Automatically detects map task characteristics by monitoring
 - Scheduling map tasks to minimize overall MapReduce job execution time



1.93 times faster than the Hadoop original scheduling



**Aoki and Nukada
“The First CUDA
Programming”
Textbook**

情報処理 2

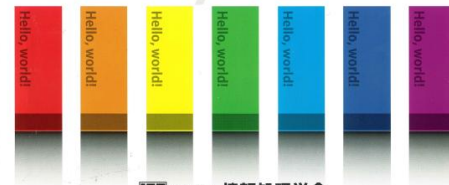
[2009] Vol.50 No.2 通巻528号



特集 アクセラレータ, 再び
—スパコン化の切り札—

解説 アウトソーシングと情報セキュリティ問題
—プリント業務のマネージド・サービスを題材として—
報告 Xen Summit Tokyo(Asia) 2008レポート
コラム わが支部の魅力はここにあり 関西支部:関西支部大会1.5倍の研究発表で支部活動の活性化

Hello, world!



社団法人 情報処理学会
Information Processing Society of Japan
<http://www.ipjs.or.jp/>

**Matsuoka, Endo,
Nukada, Aoki et. al.
Special Issue IPSJ
Magazine 2009,
“Accelerators,
Reborn”**



**TSUBAME E-Science
Journal
(Vol.2 at SC10)**



GPU Computing Consortium

**Established Sep. 1, 2009 in GSIC
607 Members as of Sep. 2010
8 CUDA Lectures (Hands-on)
1 International Workshop
1 Symposium (Oct., 2010)
3 Seminars by World-famous**