



NVIDIA® TESLA™ GPU COMPUTING SOLUTION FOR DATA CENTERS



THE WORLD'S FIRST MASS MARKET PARALLEL PROCESSORS

TESLA™ S2050 / S2070

The Tesla S2050/S2070 1U Computing Systems are designed from the ground up for high performance computing. Based on the next generation NVIDIA CUDA™ GPU architecture codenamed “Fermi”, it supports many “must have” features for technical and enterprise computing. These include ECC memory for uncompromised accuracy and scalability, support for C++ and 8X the double precision performance compared Tesla 10-series GPU computing products. When compared to the latest quad-core CPU, Tesla 20-series GPU computing systems deliver equivalent performance at 1/20th the power consumption and 1/10th the cost.

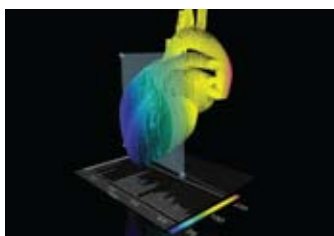
Designed with four latest-generation Tesla computing processors in a standard 1U chassis, the Tesla S2050/S2070 computing systems scales to solve the world's most important computing challenges – more quickly and accurately.



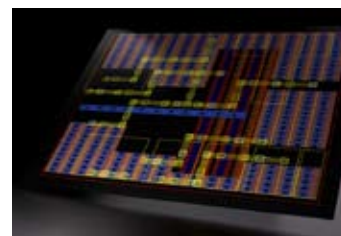
COMPUTATIONAL FLUID DYNAMICS



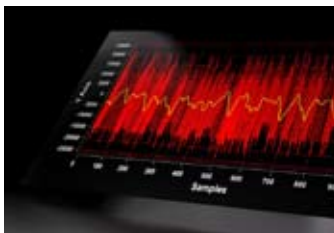
MOLECULAR DYNAMICS



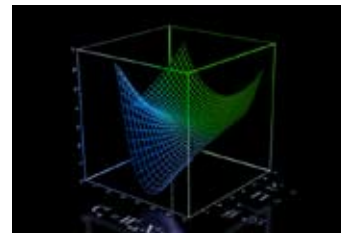
LIFE SCIENCES



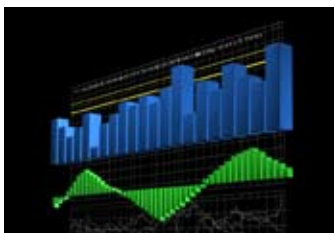
ELECTRONIC DESIGN AUTOMATION



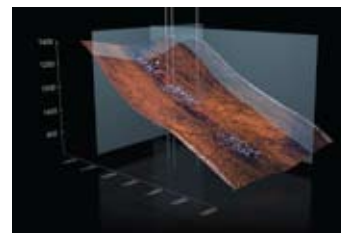
SIGNAL PROCESSING



NUMERICS



COMPUTATIONAL FINANCE



SEISMIC EXPLORATION

FEATURES AND BENEFITS

GPUs POWERED BY THE MASSIVELY PARALLEL CUDA ARCHITECTURE	Delivers cluster performance at 1/20th the power and 1/10th the cost of CPU-only systems based on the latest quad core CPUs.
IEEE 754 SINGLE AND DOUBLE PRECISION FLOATING POINT UNITS	Delivers up to 2.5 Teraflops of double precision performance in 1U form factor for faster and more accurate results.
ECC SUPPORT	Meets a critical requirement for datacenters and supercomputing centers deploying GPUs on a large scale with uncompromised computing accuracy and reliability. Offers protection of data in memory to enhance data integrity and reliability for applications. Register files, L1/L2 caches, shared memory, and DRAM all are ECC protected.
UP TO 6GB OF GDDR5 MEMORY PER GPU	Allows faster access to larger data sets.
NVIDIA® PARALLEL DATACACHE™	Accelerates algorithms such as physics solvers, ray-tracing, and sparse matrix multiplication where data addresses are not known beforehand.
NVIDIA® GIGATHREAD™ ENGINE	Maximizes the throughput by faster context switching, concurrent kernel execution, and improved thread block scheduling.
ASYNCHRONOUS TRANSFER	Turbocharges system performance by executing data transfers, even when the computing cores are busy.
SYSTEM MONITORING FEATURES	Simplifies management and monitoring post-installation. Remote capabilities as well as status lights on the front and rear of the unit ensures IT staff can see the status whether they are on the either side of the rack.
HIGH SPEED , PCIe GEN 2.0 DATA TRANSFER	Maximizes bandwidth between the host system and the Tesla processors. Enables Tesla systems to work with virtually any PCIe-compliant host system with an open PCIe slot (x8 or x16).

SPECIFICATIONS

FORM FACTOR	1.71" high × 17.425" wide × 28.5" depth
# OF TESLA GPUs	4
DOUBLE PRECISION FLOATING POINT PERFORMANCE (PEAK)	2.1 TFlops - 2.5 TFlops
TOTAL DEDICATED MEMORY TESLA S2050 TESLA S2070	12GB GDDR5 (4 × 3GB*) 24GB GDDR5 (4 × 6GB*)
POWER CONSUMPTION	900W (typical)
SYSTEM INTERFACE	PCIe x16 Gen2
SOFTWARE DEVELOPMENT TOOLS	CUDA C/C++/Fortran, OpenCL, DirectCompute Toolkits

* With ECC enabled, memory available to the user will be 2.625GB per GPU for S2050 and 5.25GB per GPU for S2070

HOST REQUIREMENTS FOR TESLA 1U SYSTEMS

- > Host system with PCIe x8 or PCIe x16
- > PCIe Gen2 for best results
- > Linux (64-bit and 32-bit)
 - > Red Hat Enterprise Linux 4 and 5
 - > SUSE 10.3
- > Windows Server 2003 and 2008

SOFTWARE DEVELOPMENT TOOLS

- > C++, complementing existing support for C, Fortran, Java, Python, OpenCL
- > Standard numerical libraries for FFT (Fast Fourier Transform), BLAS (Basic Linear Algebra Subroutines), and CuDPP (CUDA)
- > The world's first fully integrated heterogeneous computing application development environment, code named "Nexus", within Microsoft Visual Studio

To learn more about NVIDIA Tesla, go to www.nvidia.com/tesla