



NVIDIA® TESLA® P100 GPU ACCELERATOR

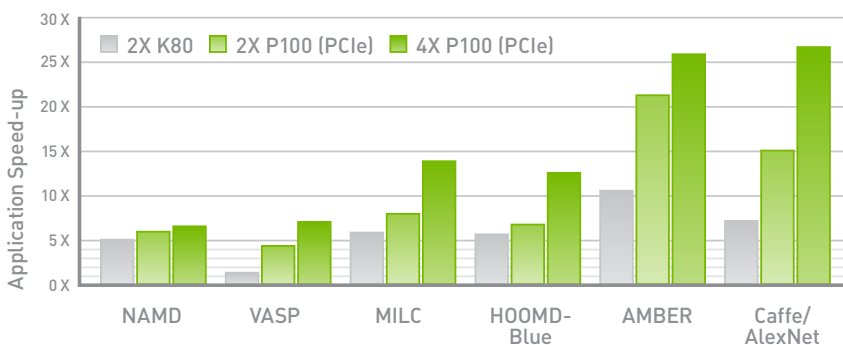
World's most advanced data center accelerator for PCIe-based servers

HPC data centers need to support the ever-growing demands of scientists and researchers while staying within a tight budget. The old approach of deploying lots of commodity compute nodes requires huge interconnect overhead that substantially increases costs without proportionally increasing performance.

NVIDIA Tesla P100 GPU accelerators are the most advanced ever built, powered by the breakthrough NVIDIA Pascal™ architecture and designed to boost throughput and save money for HPC and hyperscale data centers. The newest addition to this family, Tesla P100 for PCIe enables a single node to replace half a rack of commodity CPU nodes by delivering lightning-fast performance in a broad range of HPC applications.

MASSIVE LEAP IN PERFORMANCE

NVIDIA Tesla P100 for PCIe Performance



Dual CPU server, Intel E5-2698 v3 @ 2.3 GHz, 256 GB System Memory, Pre-Production Tesla P100



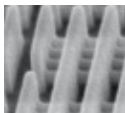
SPECIFICATIONS

GPU Architecture	NVIDIA Pascal
NVIDIA CUDA® Cores	3584
Double-Precision Performance	4.7 TeraFLOPS
Single-Precision Performance	9.3 TeraFLOPS
Half-Precision Performance	18.7 TeraFLOPS
GPU Memory	16GB CoWoS HBM2 at 732 GB/s or 12GB CoWoS HBM2 at 549 GB/s
System Interface	PCIe Gen3
Max Power Consumption	250 W
ECC	Yes
Thermal Solution	Passive
Form Factor	PCIe Full Height/Length
Compute APIs	CUDA, DirectCompute, OpenCL™, OpenACC

TeraFLOPS measurements with NVIDIA GPU Boost™ technology

A GIANT LEAP IN PERFORMANCE

Tesla P100 for PCIe is reimagined from silicon to software, crafted with innovation at every level. Each groundbreaking technology delivers a dramatic jump in performance to substantially boost the data center throughput.



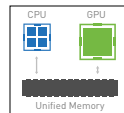
PASCAL ARCHITECTURE

More than 18.7 TeraFLOPS of FP16, 4.7 TeraFLOPS of double-precision, and 9.3 TeraFLOPS of single-precision performance powers new possibilities in deep learning and HPC workloads.



COWOS HBM2

Compute and data are integrated on the same package using Chip-on-Wafer-on-Substrate with HBM2 technology for 3X memory performance over the previous-generation architecture.



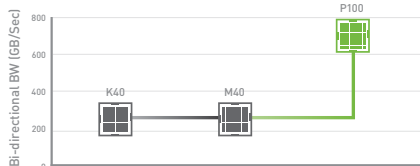
PAGE MIGRATION ENGINE

Simpler programming and computing performance tuning means that applications can now scale beyond the GPU's physical memory size to virtually limitless levels.

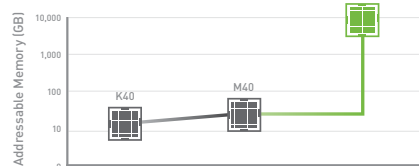
Exponential HPC and hyperscale performance



3X memory boost



Virtually limitless memory scalability



To learn more about the Tesla P100 for PCIe visit www.nvidia.com/tesla

© 2016 NVIDIA Corporation. All rights reserved. NVIDIA, the NVIDIA logo, Tesla, NVIDIA GPU Boost, CUDA, and NVIDIA Pascal are trademarks and/or registered trademarks of NVIDIA Corporation in the U.S. and other countries. OpenCL is a trademark of Apple Inc. used under license to the Khronos Group Inc. All other trademarks and copyrights are the property of their respective owners. OCT16

