



NVIDIA TESLA P6 GPU ACCELERATOR OPTIMIZED FOR BLADE SERVERS

GET MAXIMUM PERFORMANCE FOR ANY WORKLOAD

Organizations and enterprises of all sizes are entering a new era of computing. An era where deep learning, visualization, and virtualization sit side-by-side in the data center and are used by scientists, researchers, designers, renderers, engineers, and knowledge workers alike. GPUs powered by the revolutionary NVIDIA Pascal™ architecture provide the computational engine for this new era, enabling amazing user experiences by accelerating both compute and graphics virtualization at scale.

The NVIDIA® Tesla® P6 is designed for blade servers and supports multiple data center workloads, including deep learning, HPC, and graphics virtualization. It delivers higher graphics performance, improved energy efficiency, and up to twice the frame buffer compared to the NVIDIA® Tesla® M6. It also comes in an MXM form factor running at less than 90 W for high-density data centers with blade servers and converged infrastructure. Plus, it can support up to 16 concurrent users (1 GB profile).



GPU	1 NVIDIA Pascal GPU
CUDA Cores	2,048
Memory Size	16 GB GDDR5
H.264 1080p30 streams	24
Max vGPU instances	16 (1 GB Profile)
vGPU Profiles	1 GB, 2 GB, 4 GB, 8 GB, 16 GB
Form Factor	MXM (blade servers)
Power	90 W (70 W opt)
Thermal	Bare Board

VIRTUALIZE ANY WORKLOAD, ANYWHERE

The Tesla P6 GPU accelerator works with NVIDIA virtual GPU software to provide the industry's highest user performance for virtualized workstations, desktops, and applications in an MXM form factor. This lets you virtualize any application—including professional graphics applications—and deliver it out to any device, anywhere. Organizations are now virtualizing high-end applications with large, complex datasets for rendering and simulations, as well as virtualizing modern business applications.

NVIDIA vGPU software shares the power of Tesla P6 GPUs across multiple virtual workstations, desktops, and apps. This means you can deliver an immersive user experience for everyone—from office workers to mobile professionals to designers—through virtual workspaces with improved management, security, and productivity.

KEY BENEFITS

Exceptional User Experience

Get the ultimate user experience for any workload. The Tesla P6 with NVIDIA® Quadro® Virtual Data Center Workstation (Quadro vDWS) software doubles your graphics performance and supports compute workloads (CUDA and OpenCL) for every vGPU, enabling professional and design engineering workflows at peak performance. Users can count on consistent performance with the new resource scheduler, which provides deterministic QoS and eliminates the problem of a “noisy neighbor.”

Optimal Management and Monitoring

Management tools give you vGPU visibility into the host or guest level, with application-level monitoring capabilities. This allows IT to intelligently design, manage, and support their end users' experiences. End-to-end management and monitoring delivers real-time insight into GPU performance. And integration with VMware vRealize Operations (vROps), Citrix Director and XenCenter put flexibility and control in the palm of your hand.

Flexible GPU Infrastructure

Support up to 2X more users (1 GB profile) on the Tesla P6 compared to the Tesla M6, for scaling high performance virtual graphics and compute. With the larger profile size - up to 2X larger GPU framebuffer than the Tesla M6 - you can support your most demanding users. The P6 provides utilization and flexibility to your NVIDIA Quadro vDWS solution in an energy efficient, blade optimized form factor helping you drive down overall TCO.

To learn more about NVIDIA Virtual GPU Software with Tesla GPUs visit www.nvidia.com/grid

© 2017 NVIDIA Corporation. All rights reserved. NVIDIA, the NVIDIA logo, Tesla, NVIDIA GRID, NVIDIA Maxwell, and CUDA are trademarks and/or registered trademarks of NVIDIA Corporation in the U.S. and other countries. Other company and product names may be trademarks of the respective companies with which they are associated. AUG17

