



TESLA DRIVER VERSION 418.87.00(LINUX)/426.00(WINDOWS)

RN-08625-418.87.00_426.00 _v01 | February 2020

Release Notes



TABLE OF CONTENTS

Chapter 1. Version Highlights.....	1
1.1. New Features.....	1
1.2. Fixed Issues.....	1
1.3. Known Issues.....	1
1.4. Virtualization.....	3
Chapter 2. Hardware and Software Support.....	5

Chapter 1.

VERSION HIGHLIGHTS

This section provides highlights of the NVIDIA Tesla 418 Driver, version 418.87.00 for Linux and 426.00 for Windows. For changes related to the 418 release of the NVIDIA display driver, review the file "NVIDIA_Changelog" available in the .run installer packages.

1.1. New Features

- ▶ Added support for CUDA 10.1 Update 2 (10.1.243). For more information on CUDA 10.1, refer to the CUDA Toolkit 10 Release [Notes](#)
- ▶ Added support for Red Hat Enterprise Linux (RHEL) 8.0

1.2. Fixed Issues

- ▶ Various security issues were addressed, for additional details on the med-high severity issues please review [NVIDIA Product Security](#) for more information.
- ▶ Local repository of driver packages (deb/rpm) for Linux distributions (e.g. **nvidia-diag-driver-local-repo-ubuntu1604-driver-version_1.0-1_amd64.deb**) no longer include diagnostic packages, consistent with the runfile (.run) installers.
- ▶ Fixed an issue in **nvidia-smi** when using the **--query-gpu** option.
- ▶ Updated output of **nvidia-smi** to indicate retired pages that will be blacklisted on subsequent driver reload (**nvidia-smi -q**).
- ▶ Fixed a crash in **nvidia_p2p_dma_map_pages** for GPUDirect when the allocation size is > 1GB.

1.3. Known Issues

GPU Performance Counters

The use of developer tools from NVIDIA that access various performance counters requires administrator privileges. See this [note](#) for more details. For example, reading

NVLink utilization metrics from `nvidia-smi` (`nvidia-smi nvlink -g 0`) would require administrator privileges.

NVML

NVML APIs may report incorrect values for NVLink counters (read/write). This issue will be fixed in a later release of the driver.

NoScanout Mode

NoScanout mode is no longer supported on NVIDIA Tesla products. If NoScanout mode was previously used, then the following line in the “screen” section of `/etc/X11/xorg.conf` should be removed to ensure that X server starts on Tesla products:

```
Option      "UseDisplayDevice" "None"
```

Tesla products now support one display of up to 4K resolution.

Unified Memory Support

Some Unified Memory APIs (for example, CPU page faults) are not supported on Windows in this version of the driver. Review the CUDA Programming Guide on the system requirements for Unified Memory

CUDA and unified memory is not supported when used with Linux power management states S3/S4.

IMPU FRU for Volta GPUs

The driver does not support the IPMI FRU multi-record information structure for NVLink. See the Design Guide for Tesla P100 and Tesla V100-SXM2 for more information.

Video Memory Support

For Windows 7 64-bit, this driver recognizes up to the total available video memory on Tesla cards for Direct3D and OpenGL applications.

For Windows 7 32-bit, this driver recognizes only up to 4 GB of video memory on Tesla cards for DirectX, OpenGL, and CUDA applications.

Experimental OpenCL Features

Select features in OpenCL 2.0 are available in the driver for evaluation purposes only.

The following are the features as well as a description of known issues with these features in the driver:

Device side enqueue

- The current implementation is limited to 64-bit platforms only.

- ▶ OpenCL 2.0 allows kernels to be enqueued with `global_work_size` larger than the compute capability of the NVIDIA GPU. The current implementation supports only combinations of `global_work_size` and `local_work_size` that are within the compute capability of the NVIDIA GPU. The maximum supported CUDA grid and block size of NVIDIA GPUs is available at <http://docs.nvidia.com/cuda/cuda-c-programming-guide/index.html#computecapabilities>. For a given grid dimension, the `global_work_size` can be determined by CUDA grid size x CUDA block size.
- ▶ For executing kernels (whether from the host or the device), OpenCL 2.0 supports non-uniform ND-ranges where `global_work_size` does not need to be divisible by the `local_work_size`. This capability is not yet supported in the NVIDIA driver, and therefore not supported for device side kernel enqueues.

Shared virtual memory

- ▶ The current implementation of shared virtual memory is limited to 64-bit platforms only.

1.4. Virtualization

To make use of GPU passthrough with virtual machines running Windows and Linux, the hardware platform must support the following features:

- ▶ A CPU with hardware-assisted instruction set virtualization: Intel VT-x or AMD-V.
- ▶ Platform support for I/O DMA remapping.
- ▶ On Intel platforms the DMA remapper technology is called Intel VT-d.
- ▶ On AMD platforms it is called AMD IOMMU.

Support for these feature varies by processor family, product, and system, and should be verified at the manufacturer's website.

Supported Hypervisors

The following hypervisors are supported:

Hypervisor	Notes
Citrix XenServer	Version 6.0 and later
VMware vSphere (ESX / ESXi)	Version 5.1 and later.
Red Hat KVM	Red Hat Enterprise Linux 7 with KVM
Microsoft Hyper-V	Windows Server 2016 Hyper-V Generation 2

Tesla products now support one display of up to 4K resolution.

Supported Graphics Cards

The following GPUs are supported for device passthrough:

GPU Family	Boards Supported
Turing	Tesla: T4
Volta	Tesla: V100
Pascal	Tesla: P100, P40, P4, P6
Maxwell	Tesla: M60, M40, M6, M4
Kepler	Tesla: K80, K520

Chapter 2.

HARDWARE AND SOFTWARE SUPPORT

Support for these feature varies by processor family, product, and system, and should be verified at the manufacturer's website.

Supported Operating Systems

The Release 418 driver is supported on the following operating systems:

- ▶ Windows 64-bit operating systems:
 - ▶ Microsoft Windows® Server 2019
 - ▶ Microsoft Windows® Server 2016
 - ▶ Microsoft Windows® Server 2012 R2
 - ▶ Microsoft Windows® 10
 - ▶ Microsoft Windows® 8.1 (**Not supported starting with Volta**)
 - ▶ Microsoft Windows® 7 (**Not supported starting with Volta**)
- ▶ Linux 64-bit distributions:
 - ▶ Fedora 29
 - ▶ OpenSUSE Leap 15.0
 - ▶ Red Hat Enterprise Linux 6.10 (**Deprecated**)
 - ▶ Red Hat Enterprise Linux 7.6
 - ▶ Red Hat Enterprise Linux 8.0
 - ▶ SUSE Linux Enterprise Server 15.0
 - ▶ SUSE Linux Enterprise Server 12.4 (Service Pack 4)
 - ▶ Ubuntu 16.04.6 LTS
 - ▶ Ubuntu 18.04.2 LTS

API Support

This release supports the following APIs:

- ▶ NVIDIA® CUDA® 10.1 for NVIDIA® Kepler™, Maxwell™, Pascal™, Volta™ and Turing™ GPUs
- ▶ OpenGL® 4.5
- ▶ Vulkan® 1.1
- ▶ DirectX 11
- ▶ DirectX 12 (Windows 10)
- ▶ Open Computing Language (OpenCL™ software) 1.2

Supported NVIDIA Tesla GPUs

The Tesla driver package is designed for systems that have one or more Tesla products installed. This release of the Tesla driver supports CUDA C/C++ applications and libraries that rely on the CUDA C Runtime and/or CUDA Driver API.

Tesla Server Platforms	
Product	Architecture
NVIDIA HGX-2	V100 and NVSwitch

Note that Microsoft Windows® is not supported on the HGX-2 platform. HGX-2 requires additional software - for more information see the [NVIDIA® DCGM documentation](#).

Tesla T-Series Products	
Product	GPU Architecture
NVIDIA Tesla T4	Turing

Tesla V-Series Products	
Product	GPU Architecture
NVIDIA Tesla V100	Volta

Tesla P-Series Products	
Product	GPU Architecture
NVIDIA Tesla P100	Pascal
NVIDIA Tesla P40	Pascal
NVIDIA Tesla P4	Pascal
NVIDIA Tesla P6	Pascal

Tesla K-Series Products	
Product	GPU Architecture
NVIDIA Tesla K520	Kepler
NVIDIA Tesla K80	Kepler

Tesla M-Class Products	
Product	GPU Architecture
NVIDIA Tesla M60	Maxwell
NVIDIA Tesla M40 24 GB	Maxwell
NVIDIA Tesla M40	Maxwell
NVIDIA Tesla M6	Maxwell
NVIDIA Tesla M4	Maxwell

Notice

THE INFORMATION IN THIS GUIDE AND ALL OTHER INFORMATION CONTAINED IN NVIDIA DOCUMENTATION REFERENCED IN THIS GUIDE IS PROVIDED “AS IS.” NVIDIA MAKES NO WARRANTIES, EXPRESSED, IMPLIED, STATUTORY, OR OTHERWISE WITH RESPECT TO THE INFORMATION FOR THE PRODUCT, AND EXPRESSLY DISCLAIMS ALL IMPLIED WARRANTIES OF NONINFRINGEMENT, MERCHANTABILITY, AND FITNESS FOR A PARTICULAR PURPOSE. Notwithstanding any damages that customer might incur for any reason whatsoever, NVIDIA's aggregate and cumulative liability towards customer for the product described in this guide shall be limited in accordance with the NVIDIA terms and conditions of sale for the product.

THE NVIDIA PRODUCT DESCRIBED IN THIS GUIDE IS NOT FAULT TOLERANT AND IS NOT DESIGNED, MANUFACTURED OR INTENDED FOR USE IN CONNECTION WITH THE DESIGN, CONSTRUCTION, MAINTENANCE, AND/OR OPERATION OF ANY SYSTEM WHERE THE USE OR A FAILURE OF SUCH SYSTEM COULD RESULT IN A SITUATION THAT THREATENS THE SAFETY OF HUMAN LIFE OR SEVERE PHYSICAL HARM OR PROPERTY DAMAGE (INCLUDING, FOR EXAMPLE, USE IN CONNECTION WITH ANY NUCLEAR, AVIONICS, LIFE SUPPORT OR OTHER LIFE CRITICAL APPLICATION). NVIDIA EXPRESSLY DISCLAIMS ANY EXPRESS OR IMPLIED WARRANTY OF FITNESS FOR SUCH HIGH RISK USES. NVIDIA SHALL NOT BE LIABLE TO CUSTOMER OR ANY THIRD PARTY, IN WHOLE OR IN PART, FOR ANY CLAIMS OR DAMAGES ARISING FROM SUCH HIGH RISK USES.

NVIDIA makes no representation or warranty that the product described in this guide will be suitable for any specified use without further testing or modification. Testing of all parameters of each product is not necessarily performed by NVIDIA. It is customer's sole responsibility to ensure the product is suitable and fit for the application planned by customer and to do the necessary testing for the application in order to avoid a default of the application or the product. Weaknesses in customer's product designs may affect the quality and reliability of the NVIDIA product and may result in additional or different conditions and/or requirements beyond those contained in this guide. NVIDIA does not accept any liability related to any default, damage, costs or problem which may be based on or attributable to: (i) the use of the NVIDIA product in any manner that is contrary to this guide, or (ii) customer product designs.

Other than the right for customer to use the information in this guide with the product, no other license, either expressed or implied, is hereby granted by NVIDIA under this guide. Reproduction of information in this guide is permissible only if reproduction is approved by NVIDIA in writing, is reproduced without alteration, and is accompanied by all associated conditions, limitations, and notices.

Trademarks

NVIDIA and the NVIDIA logo are trademarks and/or registered trademarks of NVIDIA Corporation in the United States and other countries. Other company and product names may be trademarks of the respective companies with which they are associated.

Copyright

© 2020 NVIDIA Corporation. All rights reserved.