



What's New in Virtual GPU Software R525 for All Supported Hypervisors

Release Notes

Table of Contents

Chapter 1. Updates by Release.....	1
1.1. Updates in Release 15.2.....	1
1.2. Updates in Release 15.1.....	2
1.3. Updates in Release 15.0.....	3

Chapter 1. Updates by Release

Updates for each release in this release family of NVIDIA vGPU software may include new features, introduction of hardware and software support, and withdrawal of hardware and software support.

1.1. Updates in Release 15.2

New Features in Release 15.2

- ▶ Support for authenticated local proxy servers by licensed clients of a Cloud License Service (CLS) instance
- ▶ Security updates - see *Security Bulletin: NVIDIA GPU Display Driver - March 2023*, which is posted shortly after the release date of this software and is listed on the [NVIDIA Product Security](#) page
- ▶ Miscellaneous bug fixes

Hardware and Software Support Introduced in Release 15.2

- ▶ Support for the for the following GPUs:
 - ▶ NVIDIA H800 PCIe 80GB
 - ▶ NVIDIA L4
- ▶ Support for Rocky Linux as a guest OS
- ▶ Support for Citrix Virtual Apps and Desktops version 7 2303

Features Deprecated in Release 15.2

The following table lists features that are deprecated in this release of NVIDIA vGPU software. Although the features remain available in this release, they might be withdrawn in a future release. In preparation for the possible removal of these features, use the preferred alternatives listed in the table.

Deprecated Feature	Preferred Alternatives
CentOS Linux as a guest OS	Rocky Linux

Deprecated Feature	Preferred Alternatives
The following CentOS Linux releases are the last releases to be supported by NVIDIA vGPU software: ▶ CentOS Linux 7.9 ▶ CentOS Linux 8 (2011)	Rocky Linux releases that are compatible with supported Red Hat Enterprise Linux releases are supported.

1.2. Updates in Release 15.1

New Features in Release 15.1

- ▶ Support for GPU System Processor (GSP) in NVIDIA vGPU deployments on GPUs based on the NVIDIA Ada Lovelace architecture
- ▶ Support for hibernation with GPU-P on Microsoft Azure Stack HCI
- ▶ Options in the NVML API and the `nvidia-smi` command for getting information about the scheduling behavior of time-sliced vGPUs
- ▶ Support for independent operation of NVIDIA CUDA Toolkit profilers on MIG-backed vGPUs on GPUs based on the NVIDIA Hopper architecture
- ▶ Support for NVIDIA Virtual Applications (vApps) on Linux OSes
- ▶ Miscellaneous bug fixes

Hardware and Software Support Introduced in Release 15.1

- ▶ Support for the for the following GPUs:
 - ▶ NVIDIA L40
 - ▶ NVIDIA RTX 6000 Ada
- ▶ Support for Windows 11 22H2 as a guest OS
- ▶ Support for Windows 10 2022 Update (22H2) as a guest OS
- ▶ Support for the following OS releases as a guest OS on the Ubuntu hypervisor:
 - ▶ Microsoft Windows 11
 - ▶ Microsoft Windows 10
- ▶ Support for Citrix Virtual Apps and Desktops version 7 2212
- ▶ Support for VMware Horizon 2212 (8.8)
- ▶ Reinstatement of support for NVIDIA Virtual GPU Management Pack for VMware vRealize Operations

Support is reinstated with the release of NVIDIA Virtual GPU Management Pack for VMware vRealize Operations 3.1. This release is compatible with the vGPU management daemon, which is based on the VMware DSDK framework.

1.3. Updates in Release 15.0

New Features in Release 15.0

- ▶ Support for NVIDIA GPUDirect[®] Storage technology on MIG-backed vGPUs
- ▶ Assignment of multiple fractional vGPUs to a single VM
A fractional vGPU is allocated only a fraction of the physical GPU's frame buffer.
- ▶ Support for NVIDIA Virtual Compute Server (vCS) in Windows guest VMs on Red Hat Enterprise Linux with KVM hypervisor
- ▶ DCH packaging of the NVIDIA vGPU software graphics driver for Windows guest OSes



Note: As a result of this change, the path to the registry key for configuring NVIDIA vGPU software licensing has changed. After an upgrade from a package that is not DCH compliant, license settings must be reconfigured in the registry key at the new path to ensure that a VM in which the driver has been upgraded can acquire a license.

- ▶ Support for a mixture of TCC and WDM operation for Windows VMs to which multiple vGPUs are assigned
- ▶ Unified memory support on the following GPUs:
 - ▶ NVIDIA A100 (all variants)
 - ▶ NVIDIA A30
- ▶ Support for non-transparent local proxy servers when NVIDIA vGPU software is served licenses by a Cloud License Service (CLS) instance
- ▶ Migration of the Virtual GPU Manager for VMware vSphere to the VMware Daemon SDK (DSDK)
- ▶ Miscellaneous bug fixes

Hardware and Software Support Introduced in Release 15.0

- ▶ Support for the following GPUs:
 - ▶ NVIDIA A100 PCIe 80GB liquid cooled
 - ▶ NVIDIA A800 PCIe 80GB
 - ▶ NVIDIA A800 PCIe 80GB liquid cooled
 - ▶ NVIDIA A800 HGX 80GB
 - ▶ NVIDIA H100 PCIe 80GB
- ▶ Support for Microsoft Azure Stack HCI
- ▶ Support for Red Hat Enterprise Linux with KVM hypervisor 9.1 and 8.7
- ▶ Support for VMware vSphere 8.0

- ▶ Support for Red Hat Enterprise Linux 9.1 and 8.7 as a guest OS
- ▶ Support for Red Hat CoreOS 4.11 as a guest OS
- ▶ Support for Citrix Virtual Apps and Desktops version 7 2209
- ▶ Support for VMware Horizon 2209 (8.7)

Feature Support Withdrawn in Release 15.0

- ▶ The legacy NVIDIA vGPU software license server is no longer supported.



Note: If you are using the legacy NVIDIA vGPU software license server to serve licenses for an earlier vGPU software release, you **must** migrate your licenses to NVIDIA License System as part of your upgrade to NVIDIA vGPU software 15.0. Otherwise, your guest VMs will **not** be able to acquire a license for NVIDIA vGPU software. For more information, refer to [Migrating Licenses from a Legacy NVIDIA vGPU Software License Server](#) in the NVIDIA License System documentation.

- ▶ All versions of Microsoft Windows Server 2016 with Hyper-V role are no longer supported as a hypervisor.
- ▶ VMware vSphere Hypervisor (ESXi) 6.7 and 6.5 are no longer supported.
- ▶ Red Hat CoreOS 4.7 is longer supported as a guest OS
- ▶ All versions of Microsoft Windows Server 2016 are no longer supported as a guest OS.
- ▶ NVIDIA Virtual GPU Management Pack for VMware vRealize Operations is no longer supported.

Support is withdrawn because the CIM provider on which the management pack depends has been replaced by a management daemon based on the VMware DSDK framework.

Notice

This document is provided for information purposes only and shall not be regarded as a warranty of a certain functionality, condition, or quality of a product. NVIDIA Corporation ("NVIDIA") makes no representations or warranties, expressed or implied, as to the accuracy or completeness of the information contained in this document and assumes no responsibility for any errors contained herein. NVIDIA shall have no liability for the consequences or use of such information or for any infringement of patents or other rights of third parties that may result from its use. This document is not a commitment to develop, release, or deliver any Material (defined below), code, or functionality.

NVIDIA reserves the right to make corrections, modifications, enhancements, improvements, and any other changes to this document, at any time without notice.

Customer should obtain the latest relevant information before placing orders and should verify that such information is current and complete.

NVIDIA products are sold subject to the NVIDIA standard terms and conditions of sale supplied at the time of order acknowledgement, unless otherwise agreed in an individual sales agreement signed by authorized representatives of NVIDIA and customer ("Terms of Sale"). NVIDIA hereby expressly objects to applying any customer general terms and conditions with regards to the purchase of the NVIDIA product referenced in this document. No contractual obligations are formed either directly or indirectly by this document.

NVIDIA products are not designed, authorized, or warranted to be suitable for use in medical, military, aircraft, space, or life support equipment, nor in applications where failure or malfunction of the NVIDIA product can reasonably be expected to result in personal injury, death, or property or environmental damage. NVIDIA accepts no liability for inclusion and/or use of NVIDIA products in such equipment or applications and therefore such inclusion and/or use is at customer's own risk.

NVIDIA makes no representation or warranty that products based on this document will be suitable for any specified use. Testing of all parameters of each product is not necessarily performed by NVIDIA. It is customer's sole responsibility to evaluate and determine the applicability of any information contained in this document, ensure the product is suitable and fit for the application planned by customer, and perform 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 document. NVIDIA accepts no 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 document or (ii) customer product designs.

No license, either expressed or implied, is granted under any NVIDIA patent right, copyright, or other NVIDIA intellectual property right under this document. Information published by NVIDIA regarding third-party products or services does not constitute a license from NVIDIA to use such products or services or a warranty or endorsement thereof. Use of such information may require a license from a third party under the patents or other intellectual property rights of the third party, or a license from NVIDIA under the patents or other intellectual property rights of NVIDIA.

Reproduction of information in this document is permissible only if approved in advance by NVIDIA in writing, reproduced without alteration and in full compliance with all applicable export laws and regulations, and accompanied by all associated conditions, limitations, and notices.

THIS DOCUMENT AND ALL NVIDIA DESIGN SPECIFICATIONS, REFERENCE BOARDS, FILES, DRAWINGS, DIAGNOSTICS, LISTS, AND OTHER DOCUMENTS (TOGETHER AND SEPARATELY, "MATERIALS") ARE BEING PROVIDED "AS IS." NVIDIA MAKES NO WARRANTIES, EXPRESSED, IMPLIED, STATUTORY, OR OTHERWISE WITH RESPECT TO THE MATERIALS, AND EXPRESSLY DISCLAIMS ALL IMPLIED WARRANTIES OF NON-INFRINGEMENT, MERCHANTABILITY, AND FITNESS FOR A PARTICULAR PURPOSE. TO THE EXTENT NOT PROHIBITED BY LAW, IN NO EVENT WILL NVIDIA BE LIABLE FOR ANY DAMAGES, INCLUDING WITHOUT LIMITATION ANY DIRECT, INDIRECT, SPECIAL, INCIDENTAL, PUNITIVE, OR CONSEQUENTIAL DAMAGES, HOWEVER CAUSED AND REGARDLESS OF THE THEORY OF LIABILITY, ARISING OUT OF ANY USE OF THIS DOCUMENT, EVEN IF NVIDIA HAS BEEN ADVISED OF THE POSSIBILITY OF SUCH DAMAGES. Notwithstanding any damages that customer might incur for any reason whatsoever, NVIDIA's aggregate and cumulative liability towards customer for the products described herein shall be limited in accordance with the Terms of Sale for the product.

VESA DisplayPort

DisplayPort and DisplayPort Compliance Logo, DisplayPort Compliance Logo for Dual-mode Sources, and DisplayPort Compliance Logo for Active Cables are trademarks owned by the Video Electronics Standards Association in the United States and other countries.

HDMI

HDMI, the HDMI logo, and High-Definition Multimedia Interface are trademarks or registered trademarks of HDMI Licensing LLC.

OpenCL

OpenCL is a trademark of Apple Inc. used under license to the Khronos Group Inc.

Trademarks

NVIDIA, the NVIDIA logo, NVIDIA GRID, NVIDIA GRID vGPU, NVIDIA Maxwell, NVIDIA Pascal, NVIDIA Turing, NVIDIA Volta, GPUDirect, Quadro, and Tesla are trademarks 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.

Copyright

© 2013-2023 NVIDIA Corporation & affiliates. All rights reserved.

