



NVIDIA H200 NVL GPU

Product Brief

Document History

PB-12128-001_v01

Version	Date	Authors	Description of Change
01	April 11, 2025	NM, DV	Initial release

Table of Contents

Chapter 1.	Overview	1
Chapter 2.	Specifications	3
2.1	Product Specifications.....	3
2.2	Environmental and Reliability Specifications	5
Chapter 3.	System Airflow Requirements	6
3.1	Airflow Direction Support	6
Chapter 4.	Product Features	7
4.1	PCI Express Interface Specifications	7
4.1.1	PCIe Support.....	7
4.1.2	Single Root I/O Virtualization Support.....	7
4.1.3	Interrupt Messaging	7
4.2	CEC Hardware Root of Trust	8
4.3	Multi-Instance GPU Support.....	8
4.4	Programmable Power	8
4.4.1	nvidia-smi	8
4.5	NVLink Bridge Support.....	9
4.5.1	NVLink Bridge	9
4.5.2	NVLink Connector Placement.....	10
4.5.3	PCIe and NVLink Topology.....	10
4.6	Form Factor	11
4.7	Power Connector.....	11
4.7.1	Power Connector Placement.....	11
4.8	Extender.....	13
Chapter 5.	Support Information.....	15
5.1	Certification.....	15
5.2	Agencies.....	15

List of Figures

Figure 1-1.	NVIDIA H200 NVL with NVLink Bridge Volumetric	2
Figure 3-1.	H200 NVL Airflow Directions.....	6
Figure 4-1.	NVLink Connector Placement – Top View	10
Figure 4-2.	NVIDIA H200 NVL PCIe Card Dimensions.....	11
Figure 4-3.	PCIe 16-Pin Power Connector.....	12
Figure 4-4.	PCIe 16-Pin Power Connector Pin Assignment	12
Figure 4-5.	Enhanced Straight Extender	14

List of Tables

Table 2-1.	Product Specifications.....	3
Table 2-2.	Memory Specifications	4
Table 2-3.	Software Specifications	4
Table 2-4.	Board Environmental and Reliability Specifications.....	5
Table 4-1.	H200 NVL NVLink Speed and Bandwidth	9
Table 4-2.	NVIDIA H200 NVL NVLink Bridge Support	10
Table 4-3.	PCIe CEM 5.1 1-Pin PCIe PSU Power Level vs. Sense Logic	13
Table 4-4.	Supported Auxiliary Power Connections	13

Chapter 1. Overview

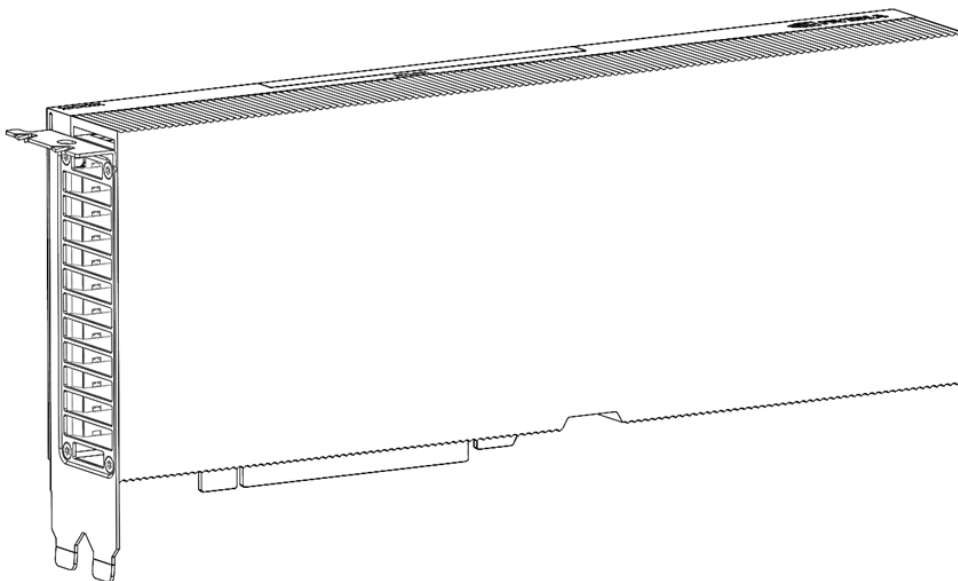
The NVIDIA® H200 NVL Tensor Core GPU is the most optimized platform for LLM inferences with its high compute density, high memory bandwidth, high energy efficiency, and unique NVIDIA® NVLink™ architecture. It also delivers unprecedented acceleration to power the world's highest-performing elastic data centers for AI, data analytics, and high-performance computing (HPC) applications. NVIDIA H200 NVL Tensor Core technology supports a broad range of math precisions, providing a single accelerator for every compute workload. The NVIDIA H200 NVL supports double precision (FP64), single-precision (FP32), half-precision (FP16), 8-bit floating point (FP8), and integer (INT8) compute tasks.

The NVIDIA H200 NVL card is a dual-slot 10.5-inch PCI Express Gen5 card based on the NVIDIA Hopper™ architecture. It uses a passive heat sink for cooling, which requires system airflow to operate the card properly within its thermal limits. The NVIDIA H200 NVL operates unconstrained up to its maximum thermal design power (TDP) level of 600 W to accelerate applications that require the fastest computational speed and highest data throughput. The NVIDIA H200 NVL debuts the world's highest PCIe card memory bandwidth -- greater than 900 gigabytes per second (GBps). This speeds time to solution for the largest models and most massive data sets.

The NVIDIA H200 NVL card features Multi-Instance GPU (MIG) capability. This can be used to partition the GPU into as many as seven hardware-isolated GPU instances, providing a unified platform that enables elastic data centers to adjust dynamically to shifting workload demands. It can allocate the right size of resources from the smallest to the biggest multi-GPU jobs. NVIDIA H200 NVL versatility means that IT managers can maximize the utility of every graphics processing unit (GPU) in their data center.

In contrast to the NVIDIA H100 and H100 NVL products, the NVIDIA H200 NVL uses a single wide NVLink bridge per card. Both 2-slot and 4-slot bridges are supported. This allows up to four NVIDIA H200 NVL PCIe cards to be connected to deliver 900 GBps bidirectional bandwidth or 14× the bandwidth of PCIe Gen5 ×16, to maximize application performance for large workloads.

Figure 1-1. NVIDIA H200 NVL with NVLink Bridge Volumetric



Chapter 2. Specifications

2.1 Product Specifications

Table 2-1 through Table 2-3 provide the product, memory, and software specifications for the NVIDIA H200 NVL card.

Table 2-1. Product Specifications

Specification	NVIDIA H200 NVL
Product SKU	P1010 SKU 230 NVPN: 900-21010-xx40-xxx
Total board power	PCIe 16-pin cable strapped for 600 W power mode: > 600 W maximum (default) > 350 W power compliance limit > 200 W minimum
Thermal solution	Passive
Mechanical form factor	Full-height, full-length (FHFL) 10.5", dual-slot
PCI Device IDs	Device ID: 0x233B Vendor ID: 0x10DE Sub-Vendor ID: 0x10DE Sub-System ID: 0x1996
GPU clocks	Base: 1,230 MHz Boost: 1,785 MHz
Performance states	P0
VBIOS	EEPROM size: 8 Mb UEFI: Supported
PCI Express interface	PCI Express Gen5 ×16; Gen5 ×8; Gen4 ×16 Lane and polarity reversal supported
Multi-Instance GPU (MIG)	Supported (seven instances)
Secure Boot (Glacier CEC1736)	Supported
Zero Power	Not supported
Power connectors and headers	One PCIe 16-pin auxiliary power connector (12VHPWR auxiliary power connector)

Specification	NVIDIA H200 NVL
Weight	Board: 1,217 grams (excluding bracket, extender, and bridge) 2-way NVLink Bridge: 49 grams per bridge 4-way NVLink Bridge: 128 grams per bridge Bracket with screws: 20 grams Enhanced straight extender: 35 grams

Table 2-2. Memory Specifications

Specification	Description
Memory clock	3,201 MHz
Memory type	HBM3e
Memory size	141 GB
Memory bus width	6016 bits
Peak memory bandwidth	4,813 GBps

Table 2-3. Software Specifications

Specification	Description ¹
SR-IOV support	Supported - 32 VF (virtual functions)
BAR address (physical function)	BAR0: 16 MiB ¹ BAR2: 256 GiB ¹ BAR4: 32 MiB ¹
BAR address (virtual function)	BAR0: 8 MiB, (256 KiB per VF) ¹ BAR1: 256 GiB, 64 bit (8 GiB per VF) ¹ BAR3: 1 GiB, 64 bit (32 MiB per VF) ¹
Message signaled interrupts	MSI-X: Supported MSI: Not supported
ARI Forwarding	Supported
Driver support	Linux: R565 TRD1 or later Windows: R565 TRD1 or later
Secure boot	Supported
CEC firmware	Version 2.0185 or later
NVFlash	Version 5.842 or later
NVIDIA® CUDA® support	x86: CUDA 12.7 or later
Virtual GPU software support	Supports vGPU 18.1 or later: NVIDIA Virtual Compute Server Edition
NVIDIA AI Enterprise	Supported with VMWare
NVIDIA certification	NVIDIA-Certified Systems™ 2.8 or later

Specification	Description ¹
PCI class code	0x03 – Display controller
PCI subclass code	0x02 – 3D controller
ECC support	Enabled
SMBus (8-bit address)	0x9E (write), 0x9F (read)
IPMI FRU EEPROM I2C address	0x50 (7-bit), 0xA0 (8 bit)
Reserved I2C addresses ²	0xAA, 0xAC, 0xA0, 0xA4
SMBus direct access	Supported
SMBPBI (SMBus Post-Box Interface)	Supported

Note:

¹The KiB, MiB, and GiB notations emphasize the “power of two” nature of the values. Thus,

> 256 KiB = 256 × 1024

> 16 MiB = 16 × 1024²

> 64 GiB = 64 × 1024³

²See Section 4.2 “CEC Hardware Root of Trust” section in this product specification.

The operator has the option to configure this power setting to be persistent across driver reloads or to revert to default power settings upon driver unload.

2.2 Environmental and Reliability Specifications

Table 2-4 provides the environmental conditions specifications for the NVIDIA H200 NVL card.

Table 2-4. Board Environmental and Reliability Specifications

Specification	Description
Ambient operating temperature	10 °C to 45 °C
Storage temperature	-40 °C to 75 °C
Operating humidity (short term) ¹	5% to 93% relative humidity
Operating humidity	5% to 85% relative humidity
Storage humidity	5% to 95% relative humidity
Mean time between failures (MTBF)	Uncontrolled environment: ² 294,332 hours at 35 °C Controlled environment: ³ 316,369 hours at 35 °C

Notes: Specifications in this table are applicable up to 6,000 feet.

¹ A period not more than 96 hours consecutive, not to exceed 15 days per year.

² Some environmental stress with limited maintenance (GF35).

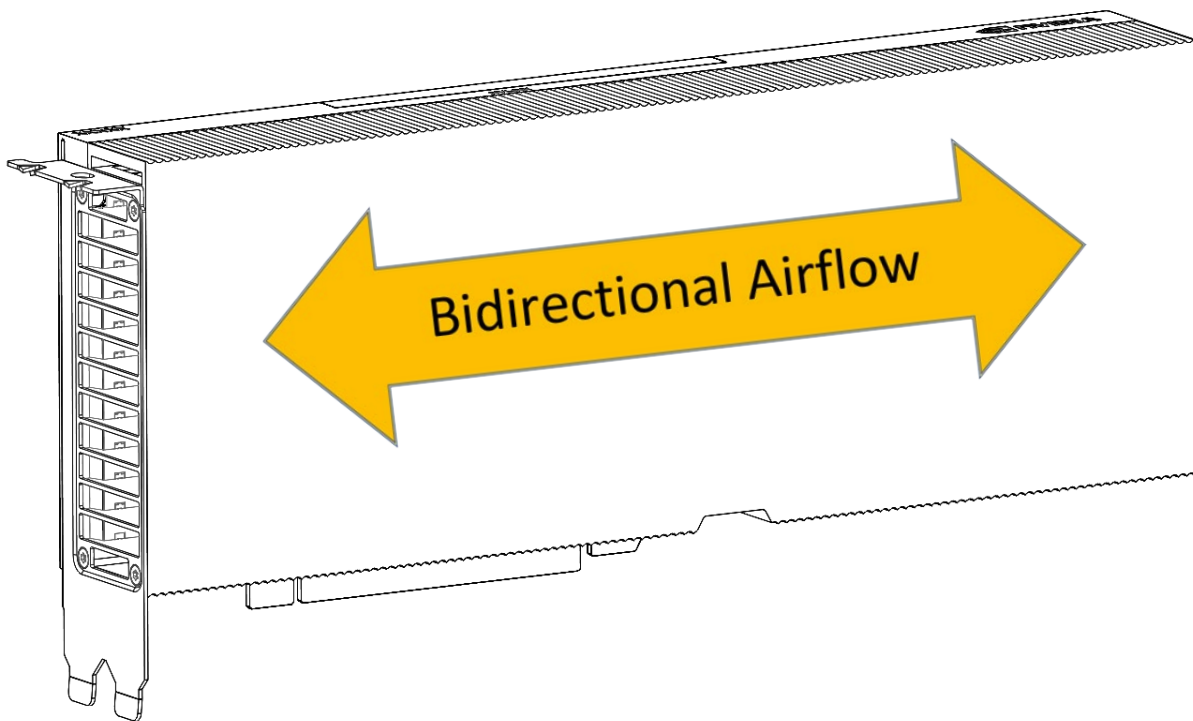
³ No environmental stress with optimum operation and maintenance (GB35).

Chapter 3. System Airflow Requirements

3.1 Airflow Direction Support

The NVIDIA H200 NVL card employs a bidirectional heat sink, which accepts airflow in either left-to-right or right-to-left directions.

Figure 3-1. H200 NVL Airflow Directions



Chapter 4. Product Features

4.1 PCI Express Interface Specifications

The following subsections describe the PCIe interface specifications for the NVIDIA H200 NVL card.

4.1.1 PCIe Support

The NVIDIA H200 NVL GPU card supports PCIe Gen5. Either a Gen5 x16, Gen5 x8, or Gen4 x16 interface should be used when connecting to the NVIDIA H200 NVL card.

4.1.2 Single Root I/O Virtualization Support

Single Root I/O Virtualization (SR-IOV) is a PCIe specification that allows a physical PCIe device to appear as multiple physical PCIe devices. Per PCIe specification, each device can have up to a maximum of 256 virtual functions (VFs). The actual number can depend on the device. SR-IOV is enabled in an NVIDIA H200 NVL card with 32 VFs supported.

For each device, SR-IOV identifies two function classes:

- > Physical functions (PFs) constitute full-featured functionality. They are fully configurable, and their configuration can control the entire device. Naturally, a PF also has the full ability to move data in and out of the device.
- > Virtual functions (VFs), which lack configuration resources. VFs exist on an underlying PF, which may support many such VFs. VFs can only move data in and out of the device. They cannot be configured and cannot be treated like a full PCIe device. The OS or hypervisor instance must be aware that they are not full PCIe devices.

The NVIDIA H200 NVL requires that SBIOS and software support in the operating system (OS) instance or hypervisor is configured to enable support for SR-IOV. The OS instance or hypervisor must be able to detect and initialize PFs and VFs.

4.1.3 Interrupt Messaging

The NVIDIA H200 NVL card only supports the MSI-X interrupt messaging protocol. The MSI interrupt protocol is not supported.

4.2 CEC Hardware Root of Trust

The NVIDIA H200 NVL provides secure boot capability using CEC. Implementing code authentication, rollback protection, and key revocation, the CEC device authenticates the contents of the GPU firmware ROM before permitting the GPU to boot from its ROM.

It also provides out-of-band (OOB) secure firmware updates, secure application processor recovery, and remote attestation.

The Hardware Root of Trust feature occupies up to two I2C addresses (in addition to the SMBus addresses). I2C addresses 0xAA and 0xAC should therefore be avoided for system use.

4.3 Multi-Instance GPU Support

The NVIDIA H200 NVL card supports Multi-Instance GPU (MIG) capability by providing up to seven GPU instances per NVIDIA H200 NVL GPU. MIG technology can partition the NVIDIA H200 NVL GPU into individual instances, each fully isolated with its own high-bandwidth memory, cache, and compute cores, enabling optimized computational resource provisioning and quality of service (QoS).

For detailed information on MIG provisioning and use, consult the *Multi-Instance GPU User's Guide*: <https://docs.nvidia.com/datacenter/tesla/mig-user-guide/index.html>

4.4 Programmable Power

The Programmable Power feature provides partners the general ability to configure the power cap of the card for system power and thermal budget or performance-per-watt reasons.

The power cap can be modified using either of these two NVIDIA tools:

- > In-band: `nvidia-smi` (power cap adjustment must be reestablished after each new driver load)
- > Out-of-band: SMBPBI (power cap adjustment remains in force across driver loads and system boots)

Power limit specifications for the NVIDIA H200 NVL are presented in Table 2-1.

4.4.1 `nvidia-smi`

`nvidia-smi` is an in-band monitoring tool provided with the NVIDIA driver and can be used to set the maximum power consumption with the driver running in persistence mode. An example command to reduce the power cap to 310 W is shown:

```
nvidia-smi -pm 1
nvidia-smi -pl 310
```

To restore the NVIDIA H200 NVL to its default TDP power consumption, either the driver module can be unloaded and reloaded, or the following command can be issued:

```
nvidia-smi -pl <600>
```

4.5 NVLink Bridge Support

NVIDIA NVLink is a high-speed point-to-point (P2P) peer transfer connection, wherein one GPU can transfer data to and receive data from other GPUs. The NVIDIA H200 NVL supports NVLink bridge connection with other adjacent NVIDIA H200 NVL cards.

The attached bridge spans either two or four PCIe slots, depending on the bridge used. Wherever an adjacent pair of the NVIDIA H200 NVL cards exists in the server, for best bridging performance and balanced bridge topology the NVIDIA H200 NVL card pair should be bridged. If there are four adjacent H200 NVL cards, the four-way NVLink bridge may be used, again using balanced bridge topology for CPU sockets.

For systems that feature multiple CPUs, bridged NVIDIA H200 NVL cards should be within the same CPU domain, that is, under the same CPU's topology. Ensuring this benefits workload application performance. There are exceptions, for example, in a system with dual CPUs wherein each CPU has a single NVIDIA H200 NVL card under it. In that case, the two NVIDIA H200 NVL cards in the system may be bridged together. See Section 4.5.3 "PCIe and NVLink Topology."

NVIDIA H200 NVL card NVLink speed and bandwidth are given in the following table.

Table 4-1. H200 NVL NVLink Speed and Bandwidth

Parameter	Value
Total NVLink bridges supported by NVIDIA H200 NVL	1
Total NVLink links supported	18
Data rate per NVIDIA H200 NVL NVLink lane (each direction)	50 gigabytes per second
Total maximum NVLink bandwidth	900 gigabytes per second

4.5.1 NVLink Bridge

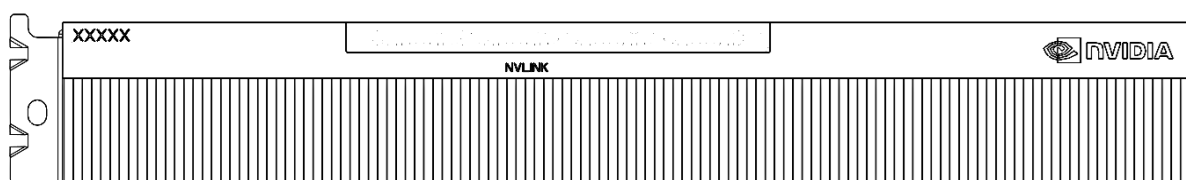
The NVLink bridge support for the NVIDIA H200 NVL cards are listed in Table 4-2.

Table 4-2. NVIDIA H200 NVL NVLink Bridge Support

NVLink Bridge	NVIDIA Part Number	Total NVLink BW
2-slot (spans two cards)	900-23945-xx00-xxx	900 gigabytes per second
4-slot (spans four cards)	900-23946-xx00-xxx	1800 gigabytes per second

4.5.2 NVLink Connector Placement

Figure 4-1 shows the connector keep-out area for NVLink bridge support on the NVIDIA H200 NVL card.

Figure 4-1. NVLink Connector Placement – Top View

Sufficient clearance must be provided both above the card's north edge and behind the backside of the card's PCB to accommodate the NVIDIA H200 NVL NVLink bridge. The clearance above the north edge should meet or exceed 2.5 mm. The backside clearance (from the rear card's rear PCB surface) should meet or exceed 2.67 mm. Consult the *NVIDIA Form Factor 5.5 Specification for Enterprise PCIe Products Specification* (NVOnline reference number 1063377) for more detailed information.

The NVLink bridge interface of the NVIDIA H200 NVL card includes a removable cap to protect the interface in non-bridged system configurations.

4.5.3 PCIe and NVLink Topology

As stated, it is strongly recommended that all NVIDIA H200 NVL cards in a set of bridge-connected cards should be within the same CPU topology domain. Bridging across CPU domains is less optimal for many workloads but is allowed if a balanced topology is maintained. Full NVLink connection topology guidance is as follows:

- > Best NVLink Topology (Recommended):
 - Bridge GPUs under the same CPU or PCIe switch
 - GPU count in a system should be in powers of two (1, 2, 4, 8, and so on)
 - Locate the same (even) number of GPUs under each CPU socket
 - Maintain a balanced configuration: same count of CPU:GPU:NIC for each grouping
 - Identical bridge slot-spans for all CPU domains
- > Good NVLink Topology:
 - Bridge GPUs under different PCIe switches but under the same CPU

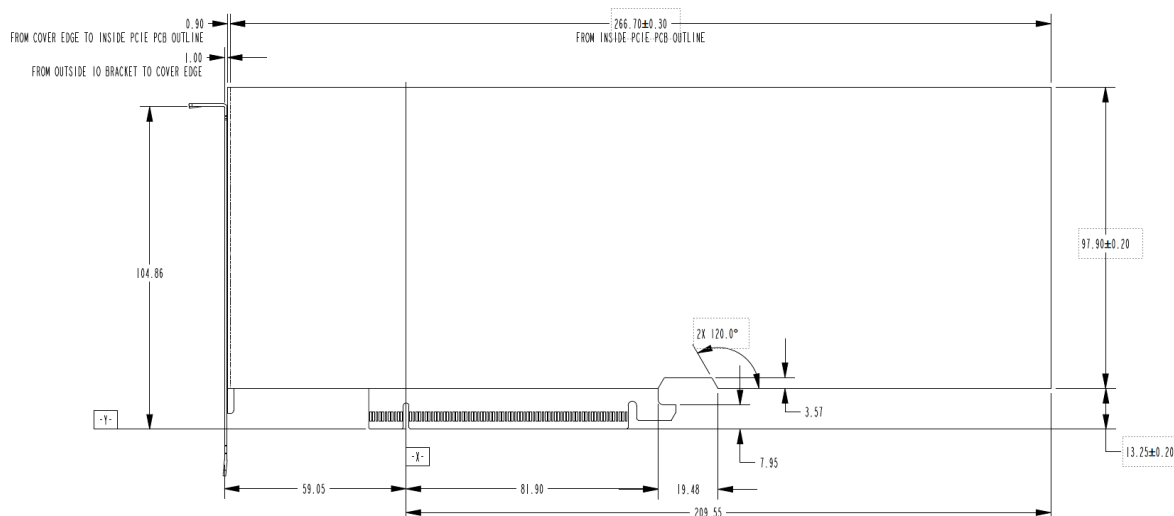
- Same number of GPUs and NICs under each CPU socket, but not powers of two
- Identical bridge slot-spans for all CPU domains
- > Allowed but Not Recommended:
 - Bridge GPUs under different CPUs
 - Odd number of GPUs under each CPU
 - Unbalanced configurations: Different ratios of CPU:GPU:NIC for each grouping

4.6 Form Factor

The NVIDIA H200 NVL card conforms to the NVIDIA Form Factor 5.5 specification for a full-height, full-length (FHFL) dual-slot PCIe card. For details, refer to the *NVIDIA Form Factor 5.5 Specification for Enterprise PCIe Products Specification* (NVOnline:1063377).

In this product specification, nominal dimensions are shown.

Figure 4-2. NVIDIA H200 NVL PCIe Card Dimensions



4.7 Power Connector

This section details the power connector for the NVIDIA H200 NVL card.

4.7.1 Power Connector Placement

The board provides a PCIe 16-pin power connector on the east edge of the board.

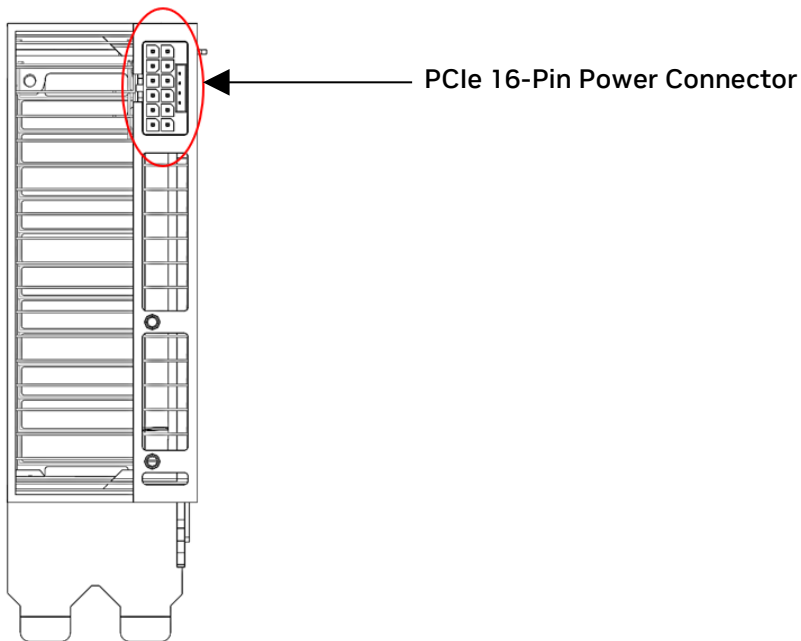
Figure 4-3. PCIe 16-Pin Power Connector

Figure 4-4 shows the pin assignments for the PCIe 16-pin power connector, per the PCIe CEM 5.1 specification.

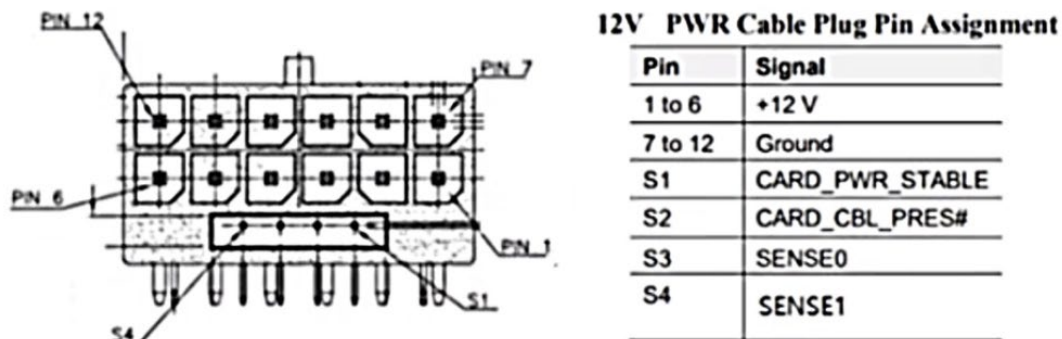
Figure 4-4. PCIe 16-Pin Power Connector Pin Assignment

Table 4-3 lists the power level options identifiable by the PCIe 16-pin power connector per the CEM 5.1 PSU, and the corresponding Sense0 and Sense1 logic. The NVIDIA card senses the Sense0 and Sense1 levels and recognizes the power available to the NVIDIA card from the power connector. If the power level identified by Sense0 and Sense1 is equal to or greater than what the NVIDIA card needs from the 16-pin connector, the NVIDIA card operates normally. If the power level identified by Sense0 and Sense1 is less than the default power cap of the NVIDIA card, the card will not boot.

The NVIDIA H200 NVL requires up to 600 W from the 16-pin auxiliary power connector. Table 4-3 shows the supported auxiliary power connector Sense pin logic and maximum supported TGP per power level. The H200 NVL requires this sense pin logic to boot, even if using the programmable TGP feature to reduce TGP. Refer to the *NVIDIA Hopper GPU Power Management for Enterprise Products* (1101509) document for details.

Table 4-3. PCIe CEM 5.1 1-Pin PCIe PSU Power Level vs. Sense Logic

Power Level	Sideband 3 (Sense0)	Sideband 4 (Sense1)	Maximum TGP
451 - 600 W	0	0	600 W
301 - 450 W	1 (float)	0	Not supported. Insufficient power
151 - 300 W	0	1 (float)	Not supported. Insufficient power
Up to 150 W	← SENSE PINS SHORTED →		Not supported. Insufficient power
No Power	1 (float)	1 (float)	Not supported. Insufficient power

Table 4-4 lists supported auxiliary power connections for the NVIDIA H200 NVL GPU card.

Table 4-4. Supported Auxiliary Power Connections

Board Connector	PSU Cable
PCIe 16 pin	PCIe 16 pin

4.8 Extender

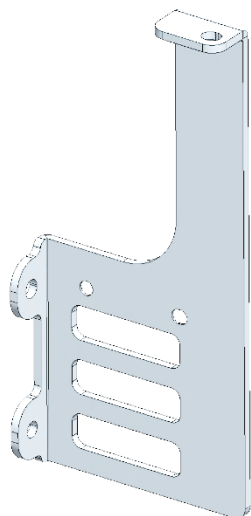
The NVIDIA H200 NVL card provides the following extender option, shown in Figure 4-5.

- > NVPN: 151-0398-000 -- Enhanced Straight Extender
 - Card + extender = 312 mm

Using a standard NVIDIA extender ensures the greatest forward compatibility with future NVIDIA product offerings.

If the standard extender will not work, OEMs may design a custom extender compliant with the applicable PCIe CEM specification or a custom attach method using the extender-mounting holes on the east edge of the PCIe card.

Figure 4-5. Enhanced Straight Extender



Chapter 5. Support Information

5.1 Certification

- > Windows Hardware Quality Lab (WHQL):
 - Windows 10, Windows 11
 - Windows Server 2019, Windows Server 2022
- > Ergonomic requirements for office work W/VDTs (ISO 9241)
- > EU Reduction of Hazardous Substances (EU RoHS)
- > Joint Industry guide (J-STD) / Registration, Evaluation, Authorization, and Restriction of Chemical Substance (EU) – (JIG / REACH)
- > Halogen Free (HF)
- > EU Waste Electrical and Electronic Equipment (WEEE)

5.2 Agencies

- > Australian Communications and Media Authority and New Zealand Radio Spectrum Management (RCM)
- > Bureau of Standards, Metrology, and Inspection (BSMI)
- > Conformité Européenne (CE)
- > Federal Communications Commission (FCC)
- > Industry Canada - Interference-Causing Equipment Standard (ICES)
- > Korean Communications Commission (KCC)
- > Underwriters Laboratories (cUL, UL)
- > Voluntary Control Council for Interference (VCCI)

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 NONINFRINGEMENT, 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.

Trademarks

NVIDIA, the NVIDIA logo, CUDA, NVIDIA-Certified Systems, NVIDIA GPU Boost, NVIDIA Hopper, and NVLink 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.

Copyright

© 2025 NVIDIA Corporation. All rights reserved.