



CUDA Quick Start Guide

Table of Contents

Chapter 1. Introduction.....	1
Chapter 2. Windows.....	2
2.1. Network Installer.....	2
2.2. Local Installer.....	4
2.3. Pip Wheels - Windows.....	6
2.4. Conda.....	8
Chapter 3. Linux.....	9
3.1. Linux x86_64.....	9
3.1.1. Redhat / CentOS.....	9
3.1.1.1. RPM Installer.....	9
3.1.1.2. Runfile Installer.....	10
3.1.2. Fedora.....	11
3.1.2.1. RPM Installer.....	11
3.1.2.2. Runfile Installer.....	12
3.1.3. SUSE Linux Enterprise Server.....	13
3.1.3.1. RPM Installer.....	13
3.1.3.2. Runfile Installer.....	14
3.1.4. OpenSUSE.....	14
3.1.4.1. RPM Installer.....	14
3.1.4.2. Runfile Installer.....	15
3.1.5. Pip Wheels - Linux.....	16
3.1.6. Conda.....	18
3.1.7. WSL.....	18
3.1.8. Ubuntu.....	18
3.1.8.1. Debian Installer.....	19
3.1.8.2. Runfile Installer.....	19
3.1.9. Debian.....	20
3.1.9.1. Debian Installer.....	20
3.1.9.2. Runfile Installer.....	21
3.2. Linux POWER8.....	22
3.2.1. Ubuntu.....	22
3.2.1.1. Debian Installer.....	22
3.2.2. Redhat / CentOS.....	23
3.2.2.1. RPM Installer.....	23
3.2.3. Conda.....	23

Chapter 1. Introduction

This guide covers the basic instructions needed to install CUDA and verify that a CUDA application can run on each supported platform.

These instructions are intended to be used on a clean installation of a supported platform. For questions which are not answered in this document, please refer to the [Windows Installation Guide](#) and [Linux Installation Guide](#).

The CUDA installation packages can be found on the [CUDA Downloads Page](#).

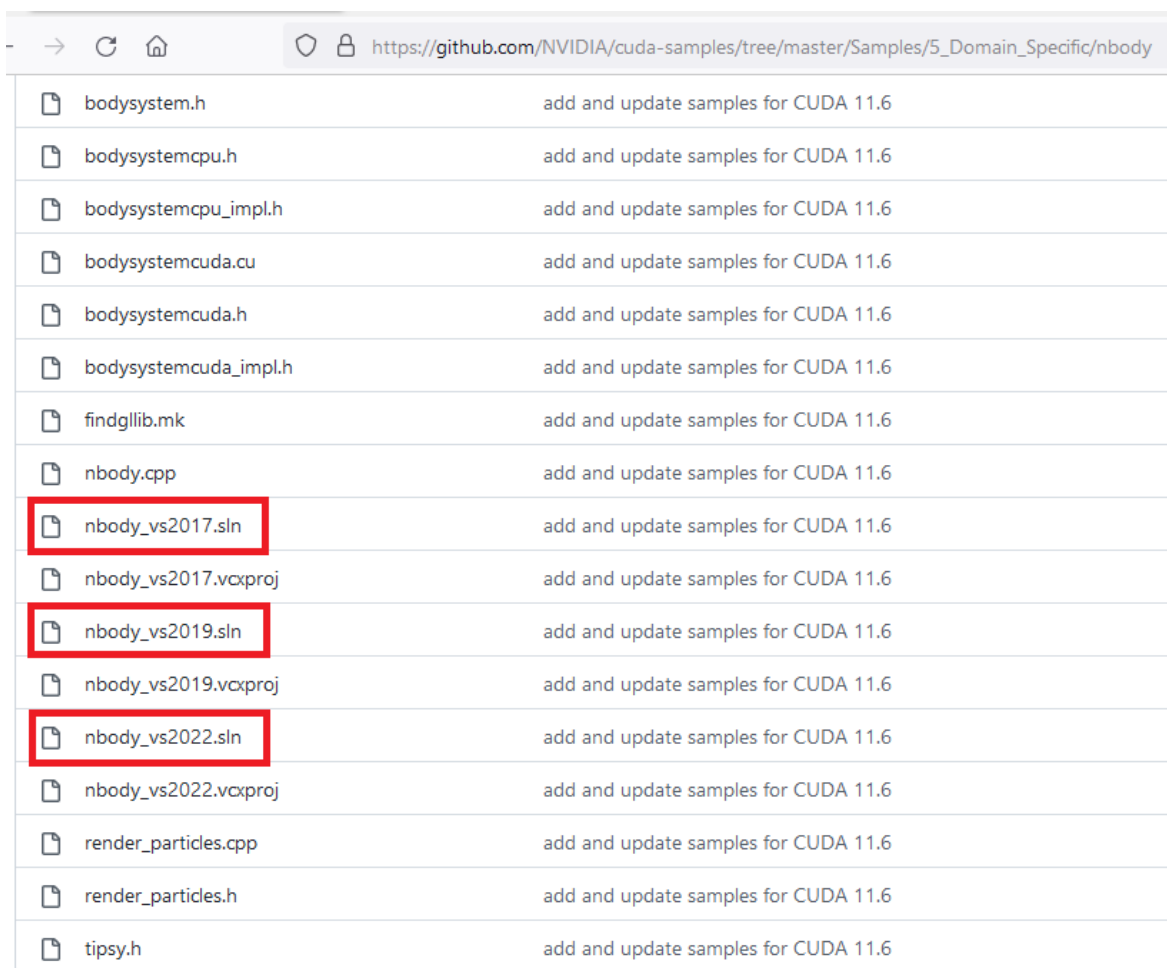
Chapter 2. Windows

When installing CUDA on Windows, you can choose between the Network Installer and the Local Installer. The Network Installer allows you to download only the files you need. The Local Installer is a stand-alone installer with a large initial download. For more details, refer to the [Windows Installation Guide](#).

2.1. Network Installer

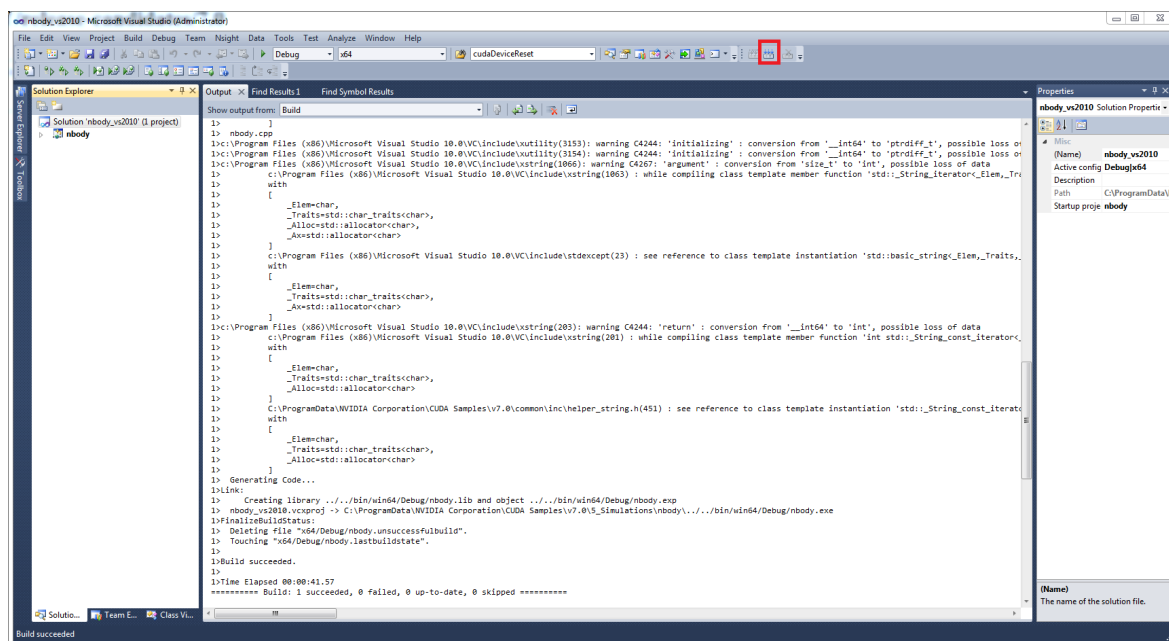
Perform the following steps to install CUDA and verify the installation.

1. Launch the downloaded installer package.
2. Read and accept the EULA.
3. Select "next" to download and install all components.
4. Once the download completes, the installation will begin automatically.
5. Once the installation completes, click "next" to acknowledge the Nsight Visual Studio Edition installation summary.
6. Click "close" to close the installer.
7. Navigate to the Samples' `nbody` directory in https://github.com/NVIDIA/cuda-samples/tree/master/Samples/5_Domain_Specific/nbody.
8. Open the `nbody` Visual Studio solution file for the version of Visual Studio you have installed, for example, `nbody_vs2019.sln`.



→ ↺ 🏠	🔒 https://github.com/NVIDIA/cuda-samples/tree/master/Samples/5_Domain_Specific/nbody
📄 bodysystem.h	add and update samples for CUDA 11.6
📄 bodysystemcpu.h	add and update samples for CUDA 11.6
📄 bodysystemcpu_impl.h	add and update samples for CUDA 11.6
📄 bodysystemcuda.cu	add and update samples for CUDA 11.6
📄 bodysystemcuda.h	add and update samples for CUDA 11.6
📄 bodysystemcuda_impl.h	add and update samples for CUDA 11.6
📄 findgllib.mk	add and update samples for CUDA 11.6
📄 nbody.cpp	add and update samples for CUDA 11.6
📄 nbody_vs2017.sln	add and update samples for CUDA 11.6
📄 nbody_vs2017.vcxproj	add and update samples for CUDA 11.6
📄 nbody_vs2019.sln	add and update samples for CUDA 11.6
📄 nbody_vs2019.vcxproj	add and update samples for CUDA 11.6
📄 nbody_vs2022.sln	add and update samples for CUDA 11.6
📄 nbody_vs2022.vcxproj	add and update samples for CUDA 11.6
📄 render_particles.cpp	add and update samples for CUDA 11.6
📄 render_particles.h	add and update samples for CUDA 11.6
📄 tipsy.h	add and update samples for CUDA 11.6

9. Open the "Build" menu within Visual Studio and click "Build Solution".



10. Navigate to the CUDA Samples' build directory and run the nbody sample.

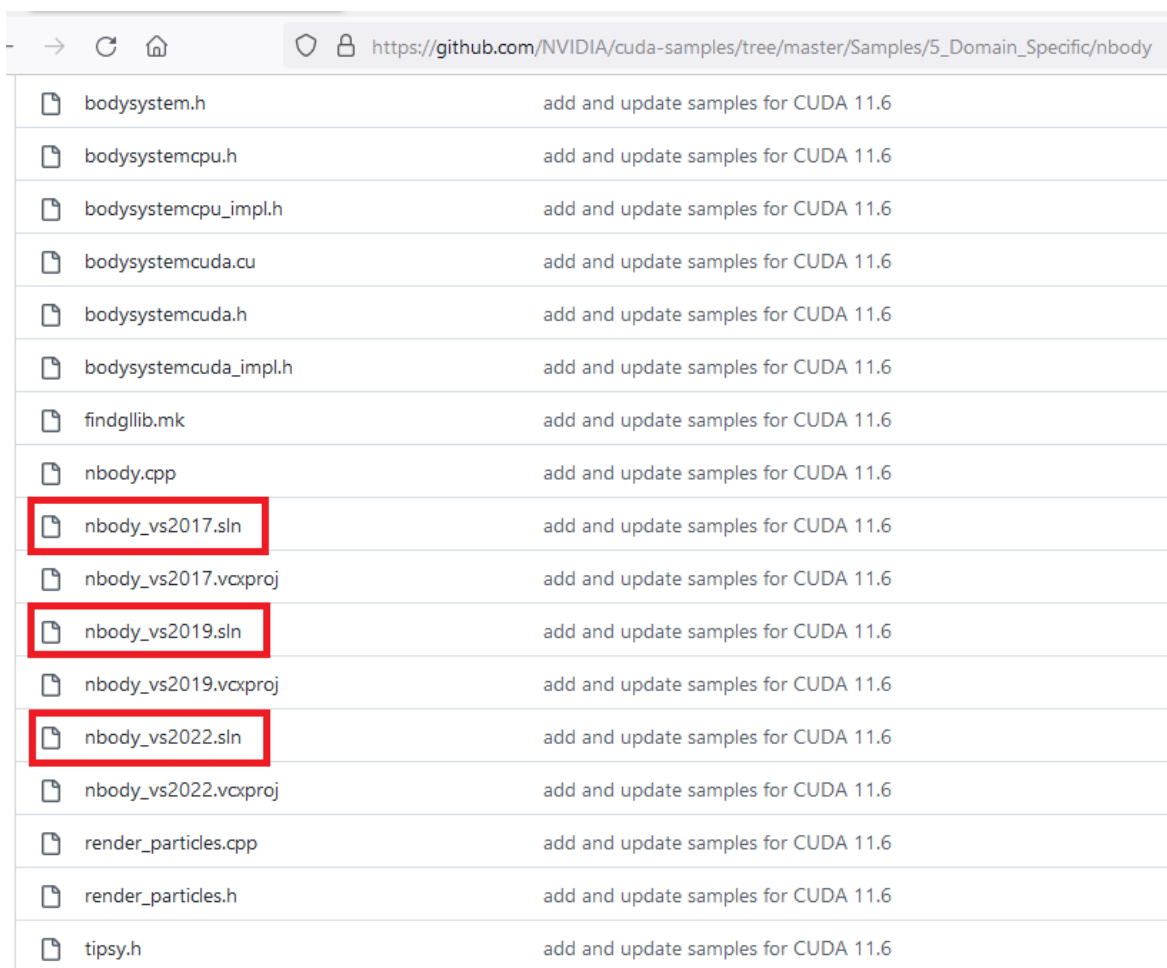


Note: Run samples by navigating to the executable's location, otherwise it will fail to locate dependent resources.

2.2. Local Installer

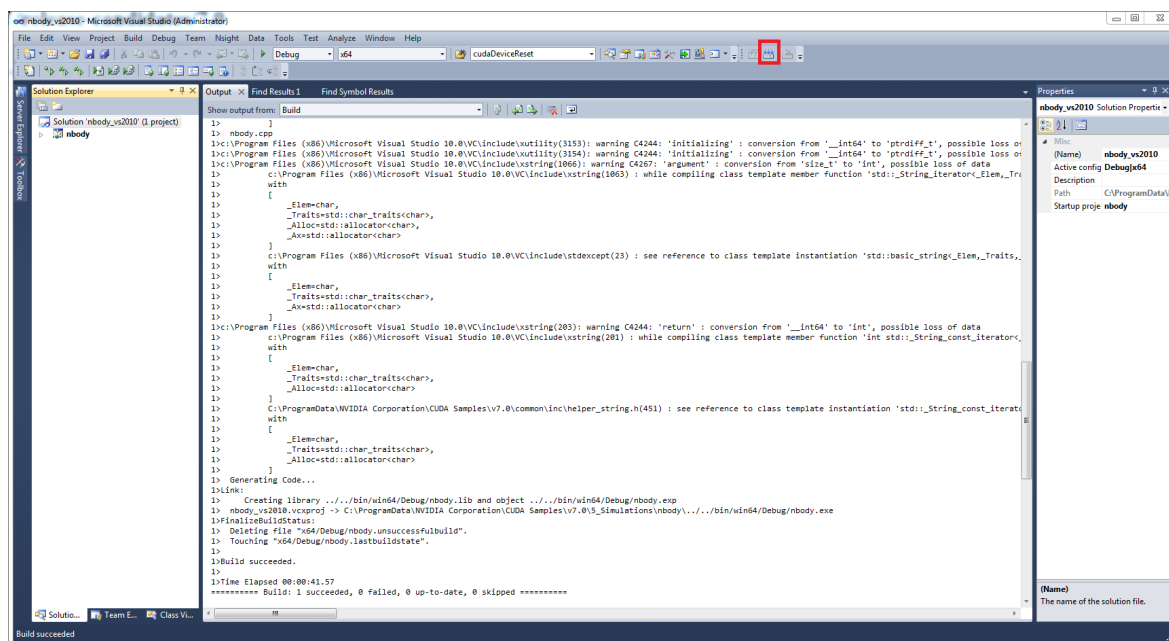
Perform the following steps to install CUDA and verify the installation.

1. Launch the downloaded installer package.
2. Read and accept the EULA.
3. Select "next" to install all components.
4. Once the installation completes, click "next" to acknowledge the Nsight Visual Studio Edition installation summary.
5. Click "close" to close the installer.
6. Navigate to the Samples' nbody directory in https://github.com/NVIDIA/cuda-samples/tree/master/Samples/5_Domain_Specific/nbody.
7. Open the nbody Visual Studio solution file for the version of Visual Studio you have installed.



→ ↺ 🏠	🔒 https://github.com/NVIDIA/cuda-samples/tree/master/Samples/5_Domain_Specific/nbody
📄 bodysystem.h	add and update samples for CUDA 11.6
📄 bodysystemcpu.h	add and update samples for CUDA 11.6
📄 bodysystemcpu_impl.h	add and update samples for CUDA 11.6
📄 bodysystemcuda.cu	add and update samples for CUDA 11.6
📄 bodysystemcuda.h	add and update samples for CUDA 11.6
📄 bodysystemcuda_impl.h	add and update samples for CUDA 11.6
📄 findgllib.mk	add and update samples for CUDA 11.6
📄 nbody.cpp	add and update samples for CUDA 11.6
📄 nbody_vs2017.sln	add and update samples for CUDA 11.6
📄 nbody_vs2017.vcxproj	add and update samples for CUDA 11.6
📄 nbody_vs2019.sln	add and update samples for CUDA 11.6
📄 nbody_vs2019.vcxproj	add and update samples for CUDA 11.6
📄 nbody_vs2022.sln	add and update samples for CUDA 11.6
📄 nbody_vs2022.vcxproj	add and update samples for CUDA 11.6
📄 render_particles.cpp	add and update samples for CUDA 11.6
📄 render_particles.h	add and update samples for CUDA 11.6
📄 tipsy.h	add and update samples for CUDA 11.6

8. Open the "Build" menu within Visual Studio and click "Build Solution".



9. Navigate to the CUDA Samples' build directory and run the nbody sample.



Note: Run samples by navigating to the executable's location, otherwise it will fail to locate dependent resources.

2.3. Pip Wheels - Windows

NVIDIA provides Python Wheels for installing CUDA through pip, primarily for using CUDA with Python. These packages are intended for runtime use and do not currently include developer tools (these can be installed separately).

Please note that with this installation method, CUDA installation environment is managed via pip and additional care must be taken to set up your host environment to use CUDA outside the pip environment.

Prerequisites

To install Wheels, you must first install the `nvidia-pyindex` package, which is required in order to set up your pip installation to fetch additional Python modules from the NVIDIA NGC PyPI repo. If your pip and setuptools Python modules are not up-to-date, then use the following command to upgrade these Python modules. If these Python modules are out-of-date then the commands which follow later in this section may fail.

```
py -m pip install --upgrade setuptools pip wheel
```

You should now be able to install the `nvidia-pyindex` module.

```
py -m pip install nvidia-pyindex
```

If your project is using a `requirements.txt` file, then you can add the following line to your `requirements.txt` file as an alternative to installing the `nvidia-pyindex` package:

```
--extra-index-url https://pypi.ngc.nvidia.com
```

Procedure

Install the CUDA runtime package:

```
py -m pip install nvidia-cuda-runtime-cu11
```

Optionally, install additional packages as listed below using the following command:

```
py -m pip install nvidia-<library>
```

Metapackages

The following metapackages will install the latest version of the named component on Windows for the indicated CUDA version. "cu11" should be read as "cuda11".

- ▶ `nvidia-cuda-runtime-cu11`
- ▶ `nvidia-cuda-cupti-cu11`
- ▶ `nvidia-cuda-nvcc-cu11`
- ▶ `nvidia-nvml-dev-cu11`
- ▶ `nvidia-cuda-nvrtc-cu11`
- ▶ `nvidia-nvtx-cu11`
- ▶ `nvidia-cuda-sanitizer-api-cu11`
- ▶ `nvidia-cublas-cu11`
- ▶ `nvidia-cufft-cu11`
- ▶ `nvidia-curand-cu11`
- ▶ `nvidia-cusolver-cu11`
- ▶ `nvidia-cusparse-cu11`
- ▶ `nvidia-npp-cu11`
- ▶ `nvidia-nvjpeg-cu11`

These metapackages install the following packages:

- ▶ `nvidia-nvml-dev-cu114`
- ▶ `nvidia-cuda-nvcc-cu114`
- ▶ `nvidia-cuda-runtime-cu114`
- ▶ `nvidia-cuda-cupti-cu114`

- ▶ nvidia-cublas-cu114
- ▶ nvidia-cuda-sanitizer-api-cu114
- ▶ nvidia-nvtx-cu114
- ▶ nvidia-cuda-nvrtc-cu114
- ▶ nvidia-npp-cu114
- ▶ nvidia-cusparse-cu114
- ▶ nvidia-cusolver-cu114
- ▶ nvidia-curand-cu114
- ▶ nvidia-cufft-cu114
- ▶ nvidia-nvjpeg-cu114

2.4. Conda

The Conda packages are available at <https://anaconda.org/nvidia>.

Installation

To perform a basic install of all CUDA Toolkit components using Conda, run the following command:

```
conda install cuda -c nvidia
```

Uninstallation

To uninstall the CUDA Toolkit using Conda, run the following command:

```
conda remove cuda
```

Chapter 3. Linux

CUDA on Linux can be installed using an RPM, Debian, Runfile, or Conda package, depending on the platform being installed on.

3.1. Linux x86_64

For development on the x86_64 architecture. In some cases, x86_64 systems may act as host platforms targeting other architectures. See the [Linux Installation Guide](#) for more details.

3.1.1. Redhat / CentOS

When installing CUDA on Redhat or CentOS, you can choose between the Runfile Installer and the RPM Installer. The Runfile Installer is only available as a Local Installer. The RPM Installer is available as both a Local Installer and a Network Installer. The Network Installer allows you to download only the files you need. The Local Installer is a stand-alone installer with a large initial download. In the case of the RPM installers, the instructions for the Local and Network variants are the same. For more details, refer to the [Linux Installation Guide](#).

3.1.1.1. RPM Installer

Perform the following steps to install CUDA and verify the installation.

1. Install EPEL to satisfy the DKMS dependency by following the instructions at [EPEL's website](#).

2. **Enable optional repos:**

On **RHEL 7 Linux** only, execute the following steps to enable optional repositories.

- ▶ **On x86_64 workstation:**

```
subscription-manager repos --enable=rhel-7-workstation-optional-rpms
```

- ▶ **On POWER9 system:**

```
subscription-manager repos --enable=rhel-7-for-power-9-optional-rpms
```

- ▶ **On x86_64 server:**

```
subscription-manager repos --enable=rhel-7-server-optional-rpms
```

3. Install the repository meta-data, clean the yum cache, and install CUDA:

```
sudo rpm --install cuda-repo-<distro>-<version>.<architecture>.rpm
sudo rpm --erase gpg-pubkey-7fa2af80*
sudo yum clean expire-cache
sudo yum install cuda
```

4. Reboot the system to load the NVIDIA drivers:

```
sudo reboot
```

5. Set up the development environment by modifying the PATH and LD_LIBRARY_PATH variables:

```
export PATH=/usr/local/cuda-11.8/bin${PATH:+:${PATH}}
export LD_LIBRARY_PATH=/usr/local/cuda-11.8/lib64\
${LD_LIBRARY_PATH:+:${LD_LIBRARY_PATH}}
```

6. Install a writable copy of the samples from <https://github.com/nvidia/cuda-samples>, then build and run the nbody sample using the Linux instructions in https://github.com/NVIDIA/cuda-samples/tree/master/Samples/5_Domain_Specific/nbody.



Note: Run samples by navigating to the executable's location, otherwise it will fail to locate dependent resources.

3.1.1.2. Runfile Installer

Perform the following steps to install CUDA and verify the installation.

1. Disable the Nouveau drivers:
 - a). Create a file at `/etc/modprobe.d/blacklist-nouveau.conf` with the following contents:

```
blacklist nouveau
options nouveau modeset=0
```

- b). Regenerate the kernel initramfs:

```
sudo dracut --force
```

2. Reboot into runlevel 3 by temporarily adding the number "3" and the word "nomodeset" to the end of the system's kernel boot parameters.
3. Run the installer silently to install with the default selections (implies acceptance of the EULA):

```
sudo sh cuda_<version>_linux.run --silent
```

4. Create an xorg.conf file to use the NVIDIA GPU for display:

```
sudo nvidia-xconfig
```

5. Reboot the system to load the graphical interface:

```
sudo reboot
```

6. Set up the development environment by modifying the PATH and LD_LIBRARY_PATH variables:

```
export PATH=/usr/local/cuda-11.8/bin${PATH:+:${PATH}}
export LD_LIBRARY_PATH=/usr/local/cuda-11.8/lib64\
${LD_LIBRARY_PATH:+:${LD_LIBRARY_PATH}}
```

7. Install a writable copy of the samples from <https://github.com/nvidia/cuda-samples>, then build and run the nbody sample using the Linux instructions in https://github.com/NVIDIA/cuda-samples/tree/master/Samples/5_Domain_Specific/nbody.



Note: Run samples by navigating to the executable's location, otherwise it will fail to locate dependent resources.

3.1.2. Fedora

When installing CUDA on Fedora, you can choose between the Runfile Installer and the RPM Installer. The Runfile Installer is only available as a Local Installer. The RPM Installer is available as both a Local Installer and a Network Installer. The Network Installer allows you to download only the files you need. The Local Installer is a stand-alone installer with a large initial download. In the case of the RPM installers, the instructions for the Local and Network variants are the same. For more details, refer to the [Linux Installation Guide](#).

3.1.2.1. RPM Installer

Perform the following steps to install CUDA and verify the installation.

1. Install the RPMFusion free repository to satisfy the Akmods dependency:

```
su -c 'dnf install --nogpgcheck http://download1.rpmfusion.org/free/fedora/
rpmfusion-free-release-$(rpm -E %fedora).noarch.rpm'
```

2. Install the repository meta-data, clean the dnf cache, and install CUDA:

```
sudo rpm --install cuda-repo-<distro>-<version>.<architecture>.rpm
sudo rpm --erase gpg-pubkey-7fa2af80*
sudo dnf clean expire-cache
sudo dnf install cuda
```

3. Reboot the system to load the NVIDIA drivers:

```
sudo reboot
```

4. Set up the development environment by modifying the PATH and LD_LIBRARY_PATH variables:

```
export PATH=/usr/local/cuda-11.8/bin${PATH:+:${PATH}}
export LD_LIBRARY_PATH=/usr/local/cuda-11.8/lib64\
${LD_LIBRARY_PATH:+:${LD_LIBRARY_PATH}}
```

5. Install a writable copy of the samples from <https://github.com/nvidia/cuda-samples>, then build and run the nbody sample using the Linux instructions in https://github.com/NVIDIA/cuda-samples/tree/master/Samples/5_Domain_Specific/nbody.



Note: Run samples by navigating to the executable's location, otherwise it will fail to locate dependent resources.

3.1.2.2. Runfile Installer

Perform the following steps to install CUDA and verify the installation.

1. Disable the Nouveau drivers:
 - a). Create a file at `/usr/lib/modprobe.d/blacklist-nouveau.conf` with the following contents:


```
blacklist nouveau
options nouveau modeset=0
```
 - b). Regenerate the kernel initramfs:


```
sudo dracut --force
```
 - c). Run the below command:


```
sudo grub2-mkconfig -o /boot/grub2/grub.cfg
```
 - d). Reboot the system:


```
sudo reboot
```
2. Reboot into runlevel 3 by temporarily adding the number "3" and the word "nomodeset" to the end of the system's kernel boot parameters.
3. Run the installer silently to install with the default selections (implies acceptance of the EULA):


```
sudo sh cuda_<version>_linux.run --silent
```
4. Create an `xorg.conf` file to use the NVIDIA GPU for display:


```
sudo nvidia-xconfig
```
5. Reboot the system to load the graphical interface.
6. Set up the development environment by modifying the `PATH` and `LD_LIBRARY_PATH` variables:

```
export PATH=/usr/local/cuda-11.8/bin${PATH:+:${PATH}}
export LD_LIBRARY_PATH=/usr/local/cuda-11.8/lib64\
    ${LD_LIBRARY_PATH:+:${LD_LIBRARY_PATH}}
```

7. Install a writable copy of the samples from <https://github.com/nvidia/cuda-samples>, then build and run the nbody sample using the Linux instructions in https://github.com/NVIDIA/cuda-samples/tree/master/Samples/5_Domain_Specific/nbody.



Note: Run samples by navigating to the executable's location, otherwise it will fail to locate dependent resources.

3.1.3. SUSE Linux Enterprise Server

When installing CUDA on SUSE Linux Enterprise Server, you can choose between the Runfile Installer and the RPM Installer. The Runfile Installer is only available as a Local Installer. The RPM Installer is available as both a Local Installer and a Network Installer. The Network Installer allows you to download only the files you need. The Local Installer is a stand-alone installer with a large initial download. In the case of the RPM installers, the instructions for the Local and Network variants are the same. For more details, refer to the [Linux Installation Guide](#).

3.1.3.1. RPM Installer

Perform the following steps to install CUDA and verify the installation.

1. Install the repository meta-data, refresh the Zypper cache, update the GPG key, and install CUDA:

```
sudo rpm --install cuda-repo-<distro>-<version>.<architecture>.rpm
sudo SUSEConnect --product PackageHub/15/x86_64
sudo zypper refresh
sudo rpm --erase gpg-pubkey-7fa2af80*
sudo dnf config-manager --add-repo https://developer.download.nvidia.com/compute/cuda/repos/$distro/$arch/cuda-$distro.repo
sudo zypper install cuda
```

2. Add the user to the video group:

```
sudo usermod -a -G video <username>
```

3. Reboot the system to load the NVIDIA drivers:

```
sudo reboot
```

4. Set up the development environment by modifying the PATH and LD_LIBRARY_PATH variables:

```
export PATH=/usr/local/cuda-11.8/bin${PATH:+:${PATH}}
export LD_LIBRARY_PATH=/usr/local/cuda-11.8/lib64\
${LD_LIBRARY_PATH:+:${LD_LIBRARY_PATH}}
```

5. Install a writable copy of the samples from <https://github.com/nvidia/cuda-samples>, then build and run the vectorAdd sample using the Linux instructions in https://github.com/NVIDIA/cuda-samples/tree/master/Samples/0_Introduction/vectorAdd.



Note: Run samples by navigating to the executable's location, otherwise it will fail to locate dependent resources.

3.1.3.2. Runfile Installer

Perform the following steps to install CUDA and verify the installation.

1. Reboot into runlevel 3 by temporarily adding the number "3" and the word "nomodeset" to the end of the system's kernel boot parameters.
2. Run the installer silently to install with the default selections (implies acceptance of the EULA):

```
sudo sh cuda_<version>_linux.run --silent
```

3. Create an xorg.conf file to use the NVIDIA GPU for display:

```
sudo nvidia-xconfig
```

4. Reboot the system to load the graphical interface:

```
sudo reboot
```

5. Set up the development environment by modifying the PATH and LD_LIBRARY_PATH variables:

```
export PATH=/usr/local/cuda-11.8/bin${PATH:+:${PATH}}
export LD_LIBRARY_PATH=/usr/local/cuda-11.8/lib64\
    ${LD_LIBRARY_PATH:+:${LD_LIBRARY_PATH}}
```

6. Install a writable copy of the samples from <https://github.com/nvidia/cuda-samples>, then build and run the vectorAdd sample using the Linux instructions in https://github.com/NVIDIA/cuda-samples/tree/master/Samples/0_Introduction/vectorAdd.



Note: Run samples by navigating to the executable's location, otherwise it will fail to locate dependent resources.

3.1.4. OpenSUSE

When installing CUDA on OpenSUSE, you can choose between the Runfile Installer and the RPM Installer. The Runfile Installer is only available as a Local Installer. The RPM Installer is available as both a Local Installer and a Network Installer. The Network Installer allows you to download only the files you need. The Local Installer is a stand-alone installer with a large initial download. In the case of the RPM installers, the instructions for the Local and Network variants are the same. For more details, refer to the [Linux Installation Guide](#).

3.1.4.1. RPM Installer

Perform the following steps to install CUDA and verify the installation.

1. Install the repository meta-data, refresh the Zypper cache, and install CUDA:

```
sudo rpm --install cuda-repo-<distro>-<version>.<architecture>.rpm
sudo rpm --erase gpg-pubkey-7fa2af80*
sudo zypper refresh
sudo zypper install cuda
```

2. Add the user to the video group:

```
sudo usermod -a -G video <username>
```

3. Reboot the system to load the NVIDIA drivers:

```
sudo reboot
```

4. Set up the development environment by modifying the PATH and LD_LIBRARY_PATH variables:

```
export PATH=/usr/local/cuda-11.8/bin${PATH:+:${PATH}}
export LD_LIBRARY_PATH=/usr/local/cuda-11.8/lib64\
${LD_LIBRARY_PATH:+:${LD_LIBRARY_PATH}}
```

5. Install a writable copy of the samples from <https://github.com/nvidia/cuda-samples>, then build and run the nbody sample using the Linux instructions in https://github.com/NVIDIA/cuda-samples/tree/master/Samples/5_Domain_Specific/nbody.



Note: Run samples by navigating to the executable's location, otherwise it will fail to locate dependent resources.

3.1.4.2. Runfile Installer

Perform the following steps to install CUDA and verify the installation.

1. Disable the Nouveau drivers:

- a). Create a file at `/etc/modprobe.d/blacklist-nouveau.conf` with the following contents:

```
blacklist nouveau
options nouveau modeset=0
```

- b). Regenerate the kernel initrd:

```
sudo /sbin/mkinitrd
```

2. Reboot into runlevel 3 by temporarily adding the number "3" and the word "nomodeset" to the end of the system's kernel boot parameters.
3. Run the installer silently to install with the default selections (implies acceptance of the EULA):

```
sudo sh cuda_<version>_linux.run --silent
```

4. Create an xorg.conf file to use the NVIDIA GPU for display:

```
sudo nvidia-xconfig
```

5. Reboot the system to load the graphical interface:

```
sudo reboot
```

6. Set up the development environment by modifying the PATH and LD_LIBRARY_PATH variables:

```
export PATH=/usr/local/cuda-11.8/bin${PATH:+:${PATH}}
export LD_LIBRARY_PATH=/usr/local/cuda-11.8/lib64\
```

```
${LD_LIBRARY_PATH:+:${LD_LIBRARY_PATH}}
```

7. Install a writable copy of the samples from <https://github.com/nvidia/cuda-samples>, then build and run the nbody sample using the Linux instructions in https://github.com/NVIDIA/cuda-samples/tree/master/Samples/5_Domain_Specific/nbody.



Note: Run samples by navigating to the executable's location, otherwise it will fail to locate dependent resources.

3.1.5. Pip Wheels - Linux

NVIDIA provides Python Wheels for installing CUDA through pip, primarily for using CUDA with Python. These packages are intended for runtime use and do not currently include developer tools (these can be installed separately).

Please note that with this installation method, CUDA installation environment is managed via pip and additional care must be taken to set up your host environment to use CUDA outside the pip environment.

Prerequisites

To install Wheels, you must first install the `nvidia-pyindex` package, which is required in order to set up your pip installation to fetch additional Python modules from the NVIDIA NGC PyPI repo. If your pip and setuptools Python modules are not up-to-date, then use the following command to upgrade these Python modules. If these Python modules are out-of-date then the commands which follow later in this section may fail.

```
python3 -m pip install --upgrade setuptools pip wheel
```

You should now be able to install the `nvidia-pyindex` module.

```
python3 -m pip install nvidia-pyindex
```

If your project is using a `requirements.txt` file, then you can add the following line to your `requirements.txt` file as an alternative to installing the `nvidia-pyindex` package:

```
--extra-index-url https://pypi.ngc.nvidia.com
```

Procedure

Install the CUDA runtime package:

```
python3 -m pip install nvidia-cuda-runtime-cu11
```

Optionally, install additional packages as listed below using the following command:

```
python3 -m pip install nvidia-<library>
```

Metapackages

The following metapackages will install the latest version of the named component on Linux for the indicated CUDA version. "cu11" should be read as "cuda11".

- `nvidia-cuda-runtime-cu11`

- ▶ `nvidia-cuda-cupti-cu11`
- ▶ `nvidia-cuda-nvcc-cu11`
- ▶ `nvidia-nvml-dev-cu11`
- ▶ `nvidia-cuda-nvrtc-cu11`
- ▶ `nvidia-nvtx-cu11`
- ▶ `nvidia-cuda-sanitizer-api-cu11`
- ▶ `nvidia-cublas-cu11`
- ▶ `nvidia-cufft-cu11`
- ▶ `nvidia-curand-cu11`
- ▶ `nvidia-cusolver-cu11`
- ▶ `nvidia-cuspars-cu11`
- ▶ `nvidia-npp-cu11`
- ▶ `nvidia-nvjpeg-cu11`

These metapackages install the following packages:

- ▶ `nvidia-nvml-dev-cu114`
- ▶ `nvidia-cuda-nvcc-cu114`
- ▶ `nvidia-cuda-runtime-cu114`
- ▶ `nvidia-cuda-cupti-cu114`
- ▶ `nvidia-cublas-cu114`
- ▶ `nvidia-cuda-sanitizer-api-cu114`
- ▶ `nvidia-nvtx-cu114`
- ▶ `nvidia-cuda-nvrtc-cu114`
- ▶ `nvidia-npp-cu114`
- ▶ `nvidia-cuspars-cu114`
- ▶ `nvidia-cusolver-cu114`
- ▶ `nvidia-curand-cu114`
- ▶ `nvidia-cufft-cu114`
- ▶ `nvidia-nvjpeg-cu114`

3.1.6. Conda

The Conda packages are available at <https://anaconda.org/nvidia>.

Installation

To perform a basic install of all CUDA Toolkit components using Conda, run the following command:

```
conda install cuda -c nvidia
```

Uninstallation

To uninstall the CUDA Toolkit using Conda, run the following command:

```
conda remove cuda
```

3.1.7. WSL

These instructions must be used if you are installing in a WSL environment. Do not use the Ubuntu instructions in this case.

1. Install repository meta-data

```
sudo dpkg -i cuda-repo-<distro>_<version>_<architecture>.deb
```

2. Update the CUDA public GPG key

```
sudo apt-key del 7fa2af80
```

When installing using the local repo:

```
sudo cp /var/cuda-repo-ubuntu2004-11-7-local/cuda-*-keyring.gpg /usr/share/
keyrings/
```

When installing using the network repo:

```
wget https://developer.download.nvidia.com/compute/cuda/repos/wsl-ubuntu/x86_64/
cuda-keyring_1.0-1_all.deb
sudo dpkg -i cuda-keyring_1.0-1_all.deb
```

Pin file to prioritize CUDA repository:

```
wget https://developer.download.nvidia.com/compute/cuda/repos/<distro>/
<architecture>/cuda-<distro>.pin
sudo mv cuda-<distro>.pin /etc/apt/preferences.d/cuda-repository-pin-600
```

3. Update the Apt repository cache and install CUDA

```
sudo apt-get update
sudo apt-get install cuda
```

3.1.8. Ubuntu

When installing CUDA on Ubuntu, you can choose between the Runfile Installer and the Debian Installer. The Runfile Installer is only available as a Local Installer. The Debian Installer is

available as both a Local Installer and a Network Installer. The Network Installer allows you to download only the files you need. The Local Installer is a stand-alone installer with a large initial download. In the case of the Debian installers, the instructions for the Local and Network variants are the same. For more details, refer to the [Linux Installation Guide](#).

3.1.8.1. Debian Installer

Perform the following steps to install CUDA and verify the installation.

1. Install the repository meta-data, update the GPG key, update the apt-get cache, and install CUDA:

```
sudo dpkg --install cuda-repo-<distro>-<version>.<architecture>.deb
sudo apt-key del 7fa2af80
wget https://developer.download.nvidia.com/compute/cuda/repos/$distro/$arch/cuda-keyring_1.0-1_all.deb
sudo dpkg -i cuda-keyring_1.0-1_all.deb
sudo add-apt-repository contrib
sudo apt-get update
sudo apt-get -y install cuda
```

2. Reboot the system to load the NVIDIA drivers:

```
sudo reboot
```

3. Set up the development environment by modifying the PATH and LD_LIBRARY_PATH variables:

```
export PATH=/usr/local/cuda-11.8/bin${PATH:+:${PATH}}
export LD_LIBRARY_PATH=/usr/local/cuda-11.8/lib64\
${LD_LIBRARY_PATH:+:${LD_LIBRARY_PATH}}
```

4. Install a writable copy of the samples from <https://github.com/nvidia/cuda-samples>, then build and run the nbody sample using the Linux instructions in https://github.com/NVIDIA/cuda-samples/tree/master/Samples/5_Domain_Specific/nbody.



Note: Run samples by navigating to the executable's location, otherwise it will fail to locate dependent resources.

3.1.8.2. Runfile Installer

Perform the following steps to install CUDA and verify the installation.

1. Disable the Nouveau drivers:
 - a). Create a file at `/etc/modprobe.d/blacklist-nouveau.conf` with the following contents:

```
blacklist nouveau
options nouveau modeset=0
```

- b). Regenerate the kernel initramfs:

```
sudo update-initramfs -u
```

2. Reboot into runlevel 3 by temporarily adding the number "3" and the word "nomodeset" to the end of the system's kernel boot parameters.

3. Run the installer silently to install with the default selections (implies acceptance of the EULA):

```
sudo sh cuda_<version>_linux.run --silent
```

4. Create an `xorg.conf` file to use the NVIDIA GPU for display:

```
sudo nvidia-xconfig
```

5. Reboot the system to load the graphical interface:

```
sudo reboot
```

6. Set up the development environment by modifying the `PATH` and `LD_LIBRARY_PATH` variables:

```
export PATH=/usr/local/cuda-11.8/bin${PATH:+:${PATH}}
export LD_LIBRARY_PATH=/usr/local/cuda-11.8/lib64\
${LD_LIBRARY_PATH:+:${LD_LIBRARY_PATH}}
```

7. Install a writable copy of the samples from <https://github.com/nvidia/cuda-samples>, then build and run the `nbody` sample using the Linux instructions in https://github.com/NVIDIA/cuda-samples/tree/master/Samples/5_Domain_Specific/nbody.



Note: Run samples by navigating to the executable's location, otherwise it will fail to locate dependent resources.

3.1.9. Debian

When installing CUDA on Debian 10, you can choose between the Runfile Installer and the Debian Installer. The Runfile Installer is only available as a Local Installer. The Debian Installer is available as both a Local Installer and a Network Installer. The Network Installer allows you to download only the files you need. The Local Installer is a stand-alone installer with a large initial download. For more details, refer to the [Linux Installation Guide](#).

3.1.9.1. Debian Installer

Perform the following steps to install CUDA and verify the installation.

1. Install the repository meta-data, remove old GPG key, install GPG key, update the apt-get cache, and install CUDA:

```
sudo dpkg -i cuda-repo-<distro>_<version>_<architecture>.deb
sudo apt-key adv --fetch-keys https://developer.download.nvidia.com/compute/cuda/
repos/debian10/x86_64/7fa2af80.pub
sudo apt-key del 7fa2af80
wget https://developer.download.nvidia.com/compute/cuda/repos/$distro/$arch/cuda-
keyring_1.0-1_all.deb
sudo dpkg -i cuda-keyring_1.0-1_all.deb
sudo add-apt-repository contrib
sudo apt-get update
sudo apt-get -y install cuda
```

2. Reboot the system to load the NVIDIA drivers:

```
sudo reboot
```

3. Set up the development environment by modifying the PATH and LD_LIBRARY_PATH variables:

```
export PATH=/usr/local/cuda-11.8/bin${PATH:+:${PATH}}
export LD_LIBRARY_PATH=/usr/local/cuda-11.8/lib64\
${LD_LIBRARY_PATH:+:${LD_LIBRARY_PATH}}
```

4. Install a writable copy of the samples from <https://github.com/nvidia/cuda-samples>, then build and run the nbody sample using the Linux instructions in https://github.com/NVIDIA/cuda-samples/tree/master/Samples/5_Domain_Specific/nbody.



Note: Run samples by navigating to the executable's location, otherwise it will fail to locate dependent resources.

3.1.9.2. Runfile Installer

Perform the following steps to install CUDA and verify the installation.

1. Disable the Nouveau drivers:
 - a). Create a file at /etc/modprobe.d/blacklist-nouveau.conf with the following contents:

```
blacklist nouveau
options nouveau modeset=0
```

- b). Regenerate the kernel initramfs:

```
sudo update-initramfs -u
```

2. Reboot into runlevel 3 by temporarily adding the number "3" and the word "nomodeset" to the end of the system's kernel boot parameters.
3. Run the installer silently to install with the default selections (implies acceptance of the EULA):

```
sudo sh cuda_<version>_linux.run --silent
```

4. Create an xorg.conf file to use the NVIDIA GPU for display:

```
sudo nvidia-xconfig
```

5. Reboot the system to load the graphical interface:

```
sudo reboot
```

6. Set up the development environment by modifying the PATH and LD_LIBRARY_PATH variables:

```
export PATH=/usr/local/cuda-11.8/bin${PATH:+:${PATH}}
export LD_LIBRARY_PATH=/usr/local/cuda-11.8/lib64\
${LD_LIBRARY_PATH:+:${LD_LIBRARY_PATH}}
```

7. Install a writable copy of the samples from <https://github.com/nvidia/cuda-samples>, then build and run the nbody sample using the Linux instructions in https://github.com/NVIDIA/cuda-samples/tree/master/Samples/5_Domain_Specific/nbody.



Note: Run samples by navigating to the executable's location, otherwise it will fail to locate dependent resources.

3.2. Linux POWER8

For development on the POWER8 architecture.

3.2.1. Ubuntu

When installing CUDA on Ubuntu on POWER8, you must use the Debian Installer. The Debian Installer is available as both a Local Installer and a Network Installer. The Network Installer allows you to download only the files you need. The Local Installer is a stand-alone installer with a large initial download. The instructions for the Local and Network variants are the same. For more details, refer to the [Linux Installation Guide](#).

3.2.1.1. Debian Installer

Perform the following steps to install CUDA and verify the installation.

1. Install the repository meta-data, update the apt-get cache, and install CUDA:

```
sudo dpkg --install cuda-repo-<distro>-<version>.<architecture>.deb
sudo apt-key del 7fa2af80
wget https://developer.download.nvidia.com/compute/cuda/repos/$distro/$arch/cuda-keyring_1.0-1_all.deb
sudo dpkg -i cuda-keyring_1.0-1_all.deb
sudo add-apt-repository contrib
sudo apt-get update
sudo apt-get -y install cuda
```

2. Reboot the system to load the NVIDIA drivers.
3. Set up the development environment by modifying the PATH and LD_LIBRARY_PATH variables:

```
export PATH=/usr/local/cuda-11.8/bin${PATH:+:${PATH}}
export LD_LIBRARY_PATH=/usr/local/cuda-11.8/lib64\
${LD_LIBRARY_PATH:+:${LD_LIBRARY_PATH}}
```

4. Install a writable copy of the samples from <https://github.com/nvidia/cuda-samples>, then build and run the vectorAdd sample using the Linux instructions in https://github.com/NVIDIA/cuda-samples/tree/master/Samples/0_Introduction/vectorAdd.



Note: Run samples by navigating to the executable's location, otherwise it will fail to locate dependent resources.

3.2.2. Redhat / CentOS

When installing CUDA on Redhat on POWER8, you must use the RPM Installer. The RPM Installer is available as both a Local Installer and a Network Installer. The Network Installer allows you to download only the files you need. The Local Installer is a stand-alone installer with a large initial download. The instructions for the Local and Network variants are the same. For more details, refer to the [Linux Installation Guide](#).

3.2.2.1. RPM Installer

Perform the following steps to install CUDA and verify the installation.

1. Install EPEL to satisfy the DKMS dependency by following the instructions at [EPEL's website](#).
2. Install the repository meta-data, clean the yum cache, and install CUDA:

```
sudo rpm --install cuda-repo-<distro>-<version>.<architecture>.rpm
sudo yum clean expire-cache
sudo yum install cuda
```

3. Reboot the system to load the NVIDIA drivers.
4. Set up the development environment by modifying the PATH and LD_LIBRARY_PATH variables:

```
export PATH=/usr/local/cuda-11.8/bin${PATH:+:${PATH}}
export LD_LIBRARY_PATH=/usr/local/cuda-11.8/lib64\
${LD_LIBRARY_PATH:+:${LD_LIBRARY_PATH}}
```

5. Install a writable copy of the samples from <https://github.com/nvidia/cuda-samples>, then build and run the vectorAdd sample using the Linux instructions in https://github.com/NVIDIA/cuda-samples/tree/master/Samples/0_Introduction/vectorAdd.



Note: Run samples by navigating to the executable's location, otherwise it will fail to locate dependent resources.

3.2.3. Conda

The Conda packages are available at <https://anaconda.org/nvidia>.

Installation

To perform a basic install of all CUDA Toolkit components using Conda, run the following command:

```
conda install cuda -c nvidia
```

Uninstallation

To uninstall the CUDA Toolkit using Conda, run the following command:

```
conda remove cuda
```

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.

OpenCL

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

Trademarks

NVIDIA and the NVIDIA logo 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

© 2015-2022 NVIDIA Corporation & affiliates. All rights reserved.