



CUDA Demo Suite

Release 13.0

NVIDIA Corporation

Aug 01, 2025

Contents

1	Demos	3
1.1	deviceQuery	3
1.2	vectorAdd	3
1.3	bandwidthTest	3
1.4	busGrind	4
1.5	nbody	5
1.6	oceanFFT	6
1.7	randomFog	6
2	Notices	9
2.1	Notice	9
2.2	OpenCL	10
2.3	Trademarks	10

CUDA Demo Suite

The reference guide for the CUDA Demo Suite.

The CUDA Demo Suite contains pre-built applications which use CUDA. These applications demonstrate the capabilities and details of NVIDIA GPUs.

Chapter 1. Demos

Below are the demos within the demo suite.

1.1. deviceQuery

This application enumerates the properties of the CUDA devices present in the system and displays them in a human readable format.

1.2. vectorAdd

This application is a very basic demo that implements element by element vector addition.

1.3. bandwidthTest

This application provides the memcpy bandwidth of the GPU and memcpy bandwidth across PCI-e. This application is capable of measuring device to device copy bandwidth, host to device copy bandwidth for pageable and page-locked memory, and device to host copy bandwidth for pageable and page-locked memory.

Arguments:

```
Usage: bandwidthTest [OPTION]...
Test the bandwidth for device to host, host to device, and device to device transfers

Example: measure the bandwidth of device to host pinned memory copies in the range
↪ 1024 Bytes to 102400 Bytes in 1024 Byte increments
./bandwidthTest --memory=pinned --mode=range --start=1024 --end=102400 --
↪ increment=1024 --dtoh
```

Options	Explanation
-help	Display this help menu
-csv	Print results as a CSV
-device=[deviceno] all 0,1,2,...,n	Specify the device device to be used compute cumulative bandwidth on all the de- vices Specify any particular device to be used
-memory=[MEMMODE] pageable pinned	Specify which memory mode to use pageable memory non-pageable system memory
-mode=[MODE] quick range shmoo	Specify the mode to use performs a quick measurement measures a user-specified range of values performs an intense shmoo of a large range of values
-htod	Measure host to device transfers
-dtoh	Measure device to host transfers
-dtod	Measure device to device transfers
-wc	Allocate pinned memory as write-combined
-cputiming	Force CPU-based timing always
Range Mode options -start=[SIZE] -end=[SIZE] -increment=[SIZE]	Starting transfer size in bytes Ending transfer size in bytes Increment size in bytes

1.4. busGrind

Provides detailed statistics about peer-to-peer memory bandwidth amongst GPUs present in the system as well as pinned, unpinned memory bandwidth.

Arguments:

Options	Explanation
-h	print usage
-p [0,1]	enable or disable pinned memory tests (default on)
-u [0,1]	enable or disable unpinned memory tests (default off)
-e [0,1]	enable or disable p2p enabled memory tests (default off)
-d [0,1]	enable or disable p2p disabled memory tests (default off)
-a	enable all tests
-n	disable all tests

Order of parameters matters.

Examples:

```
./BusGrind -n -p 1 -e 1    Run all pinned and P2P tests
./BusGrind -n -u 1        Runs only unpinned tests
./BusGrind -a             Runs all tests (pinned, unpinned, p2p enabled, p2p
↳ disabled)
```

1.5. nbody

This demo does an efficient all-pairs simulation of a gravitational n-body simulation in CUDA. It scales the n-body simulation across multiple GPUs in a single PC if available. Adding “-numbodies=num_of_bodies” to the command line will allow users to set # of bodies for simulation. Adding “-numdevices=N” to the command line option will cause the sample to use N devices (if available) for simulation. In this mode, the position and velocity data for all bodies are read from system memory using “zero copy” rather than from device memory. For a small number of devices (4 or fewer) and a large enough number of bodies, bandwidth is not a bottleneck so we can achieve strong scaling across these devices.

Arguments:

Options	Explanation
-fullscreen	run n-body simulation in fullscreen mode
-fp64	use double precision floating point values for simulation
-hostmem	stores simulation data in host memory
-benchmark	run benchmark to measure performance
-numbodies=N	number of bodies (≥ 1) to run in simulation
-device=d	where $d=0,1,2,\dots$ for the CUDA device to use
-numdevices=i	where $i=(\text{number of CUDA devices} > 0)$ to use for simulation
-compare	compares simulation results running once on the default GPU and once on the CPU
-cpu	run n-body simulation on the CPU
-tipsy=file.bin	load a tipsy model file for simulation

1.6. oceanFFT

This is a graphical demo which simulates an ocean height field using the CUFFT library, and renders the result using OpenGL.

The following keys can be used to control the output:

Keys	Function
w	Toggle wireframe

1.7. randomFog

This is a graphical demo which does pseudo- and quasi- random numbers visualization produced by CURAND. On creation, randomFog generates 200,000 random coordinates in spherical coordinate space (radius, angle ρ , angle θ) with curand's XORWOW algorithm. The coordinates are normalized for a uniform distribution through the sphere. The X axis is drawn with blue in the negative direction and yellow positive. The Y axis is drawn with green in the negative direction and magenta positive. The Z axis is drawn with red in the negative direction and cyan positive.

The following keys can be used to control the output:

Keys	Function
s	Generate new set of random nos and display as spherical coordinates (Sphere)
e	Generate new set of random nos and display on a spherical surface (shEll)
b	Generate new set of random nos and display as cartesian coordinates (cuBe/Box)
p	Generate new set of random nos and display on a cartesian plane (Plane)
i, l, j	Rotate the negative Z-axis up, right, down and left respectively
a	Toggle auto-rotation
t	Toggle 10x zoom
z	Toggle axes display
x	Select XORWOW generator (default)
c	Select Sobol' generator
v	Select scrambled Sobol' generator
r	Reset XORWOW (i.e. reset to initial seed) and re-generate
]	Increment the number of Sobol' dimensions and regenerate
[	Reset the number of Sobol' dimensions to 1 and regenerate
►	Increment the number of displayed points by 8,000 (max. 200,000)
◄	Decrement the number of displayed points by 8,000 (down to min. 8000)
q/[ESC]	Quit the application.

Chapter 2. Notices

2.1. 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.

2.2. OpenCL

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

2.3. 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

©2016-2025, NVIDIA Corporation & affiliates. All rights reserved