

Present-day high-performance computing (HPC) and deep learning applications benefit from, and even require, cluster-scale GPU compute power. Writing CUDA® applications that can correctly and efficiently utilize GPUs across a cluster requires a distinct set of skills. In this workshop, you'll learn the tools and techniques needed to write CUDA C++ applications that can scale efficiently to clusters of NVIDIA GPUs.

You'll do this by working on code from several CUDA C++ applications in an interactive cloud environment backed by several NVIDIA GPUs. You'll gain exposure to a handful of multi-GPU programming methods, including CUDA-aware Message Passing Interface (MPI), before proceeding to the main focus of this course, **NVSHMEM™**.

NVSHMEM is a parallel programming interface based on OpenSHMEM that provides efficient and scalable communication for NVIDIA GPU clusters. NVSHMEM creates a global address space for data that spans the memory of multiple GPUs and can be accessed with fine-grained GPU-initiated operations, CPU-initiated operations, and operations on CUDA streams. NVSHMEM's asynchronous, GPU-initiated data transfers eliminate synchronization overheads between the CPU and the GPU. They also enable long-running kernels that include both communication and computation, reducing overheads that can limit an application's performance when strong scaling. NVSHMEM has been used to great effect on systems such as the Summit supercomputer located at the Oak Ridge Leadership Computing Facility (OLCF), the Lawrence Livermore National Laboratory's Sierra supercomputer, and the NVIDIA DGX™ A100.

## Learning Objectives

By participating in the workshop, you'll:

- > Learn several methods for writing multi-GPU CUDA C++ applications
- > Use a variety of multi-GPU communication patterns and understand their tradeoffs
- > Write portable, scalable CUDA code with the single-program multiple-data (SPMD) paradigm using CUDA-aware MPI and NVSHMEM
- > Improve multi-GPU SPMD code with NVSHMEM's symmetric memory model and its ability to perform GPU-initiated data transfers
- > Get practice with common multi-GPU coding paradigms like domain decomposition and halo exchanges
- > Explore scaling considerations for a variety of GPU-cluster configurations

## Overview

|                                         |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
|-----------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>Duration</b>                         | 8 hours                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                       |
| <b>Price</b>                            | <a href="#">Contact us for pricing</a>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        |
| <b>Prerequisites</b>                    | <ul style="list-style-type: none"><li>&gt; Intermediate experience writing CUDA C/C++ applications</li></ul> <p>Suggested materials to satisfy the prerequisites:</p> <ul style="list-style-type: none"><li>&gt; <a href="#">Fundamentals of Accelerated Computing with CUDA C/C++</a></li><li>&gt; <a href="#">Accelerating CUDA C++ Applications with Multiple GPUs</a></li><li>&gt; <a href="#">Accelerating CUDA C++ Applications with Concurrent Streams</a></li><li>&gt; <a href="#">Scaling Workloads Across Multiple GPUs with CUDA C++</a></li></ul> |
| <b>Tools, libraries, and frameworks</b> | CUDA, MPI, NVSHMEM                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                            |

|                              |                                                                                                                                                                                                                 |
|------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>Assessment type</b>       | Skills-based coding assessment<br><br>Students must refactor a single-GPU 1D wave function solver to be GPU-cluster-ready with NVSHMEM.                                                                         |
| <b>Certificate</b>           | Upon successful completion of the assessment, participants will receive an NVIDIA DLI certificate to recognize their subject matter competency and support professional career growth.                          |
| <b>Hardware requirements</b> | Desktop or laptop computer capable of running the latest version of Chrome or Firefox. Each participant will be provided with dedicated access to a fully configured, GPU-accelerated workstation in the cloud. |
| <b>Language</b>              | English                                                                                                                                                                                                         |

## Workshop Outline

|                                                         |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                              |
|---------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>Introduction</b> (15 mins)                           | <ul style="list-style-type: none"> <li>&gt; Meet the instructor.</li> <li>&gt; Create an account at <a href="https://courses.nvidia.com/join">courses.nvidia.com/join</a></li> </ul>                                                                                                                                                                                                                                                                                                                                                         |
| <b>Multi-GPU Programming Paradigms</b><br>(120 minutes) | <p>Survey multiple techniques for programming CUDA C++ applications for multiple GPUs using a Monte Carlo approximation of a CUDA C++ program.</p> <ul style="list-style-type: none"> <li>&gt; Use CUDA to utilize multiple GPUs.</li> <li>&gt; Learn how to enable and use direct peer-to-peer memory communication.</li> <li>&gt; Write an SPMD version with CUDA-aware MPI.</li> </ul>                                                                                                                                                    |
| <b>Break</b> (60 minutes)                               |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                              |
| <b>Introduction to NVSHMEM</b><br>(120 minutes)         | <p>Learn how to write code with NVSHMEM and understand its symmetric memory model.</p> <ul style="list-style-type: none"> <li>&gt; Use NVSHMEM to write SPMD code for multiple GPUs.</li> <li>&gt; Utilize symmetric memory to let all GPUs access data on other GPUs.</li> <li>&gt; Make GPU-initiated memory transfers.</li> </ul>                                                                                                                                                                                                         |
| <b>Break</b> (15 minutes)                               |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                              |
| <b>Halo Exchanges with NVSHMEM</b><br>(90 minutes)      | <p>Practice common coding motifs like halo exchanges and domain decomposition using NVSHMEM, and work on the assessment.</p> <ul style="list-style-type: none"> <li>&gt; Write an NVSHMEM implementation of a Laplace equation Jacobi solver.</li> <li>&gt; Refactor a single GPU 1D wave equation solver with NVSHMEM.</li> <li>&gt; Complete the assessment and earn a certificate.</li> </ul>                                                                                                                                             |
| <b>Final Review</b><br>(45 minutes)                     | <ul style="list-style-type: none"> <li>&gt; Learn about application tradeoffs on GPU clusters.</li> <li>&gt; Review key learnings and answer questions.</li> <li>&gt; Complete the workshop survey.</li> </ul>                                                                                                                                                                                                                                                                                                                               |
| <b>Next Steps</b>                                       | <p>Continue learning with these DLI trainings:</p> <ul style="list-style-type: none"> <li>&gt; <b>Fundamentals of Deep Learning for Multi-GPUs</b> for those interested in cluster-scale deep learning.</li> <li>&gt; <b>High Performance Computing with Containers</b> for HPC programmers looking to improve the portability of their cluster-scale applications.</li> <li>&gt; <b>Fundamentals of Accelerated Computing with CUDA Python</b> for CUDA C++ programmers who would like to extend their CUDA knowledge to Python.</li> </ul> |

## Why Choose NVIDIA Deep Learning Institute for Hands-On Training?

- > Access workshops from anywhere with just your desktop/laptop and an internet connection. Each participant will have access to a fully configured, GPU-accelerated server in the cloud.
- > Obtain hands-on experience with the most widely used, industry-standard software, tools, and frameworks.
- > Learn how to build deep learning and accelerated computing applications for industries, such as healthcare, robotics, manufacturing, accelerated computing, and more.
- > Gain real-world experience through content designed in collaboration with industry leaders, such as the Children's Hospital of Los Angeles, Mayo Clinic, and PwC.
- > Earn an NVIDIA DLI certificate to demonstrate your subject matter competency and support your career growth. 🌟

For the latest DLI workshops and trainings, visit [www.nvidia.com/dli](https://www.nvidia.com/dli)

For questions, contact us at [nvdl@nvidia.com](mailto:nvdl@nvidia.com)

© 2021 NVIDIA Corporation. All rights reserved. NVIDIA, the NVIDIA logo, CUDA, and DGX are trademarks and/or registered trademarks of NVIDIA Corporation in the U.S. and other countries. All other trademarks and copyrights are the property of their respective owners. SEP21

