



KEPLER COMPATIBILITY GUIDE FOR CUDA APPLICATIONS

DA-06287-001_v5.5 | May 2013

Application Note



TABLE OF CONTENTS

Chapter 1. Kepler Compatibility	1
1.1 About this Document	1
1.2 Application Compatibility on Kepler	1
1.3 Verifying Kepler Compatibility for Existing Applications	2
1.3.1 Applications Using CUDA Toolkit 4.1 or Earlier	2
1.3.2 Applications Using CUDA Toolkit 5.0	2
1.4 Building Applications with Kepler Support	2
1.4.1 Applications Using CUDA Toolkit 4.1 or Earlier	3
1.4.2 Applications Using CUDA Toolkit 5.0	4
Appendix A. Revision History	5

Chapter 1.

KEPLER COMPATIBILITY

1.1 About this Document

This application note, *Kepler Compatibility Guide for CUDA Applications*, is intended to help developers ensure that their NVIDIA® CUDA™ applications will run effectively on GPUs based on the NVIDIA® Kepler Architecture. This document provides guidance to developers who are already familiar with programming in CUDA C/C++ and want to make sure that their software applications are compatible with Kepler.

1.2 Application Compatibility on Kepler

The NVIDIA CUDA C compiler, **nvcc**, can be used to generate both architecture-specific **cubin** files and forward-compatible PTX versions of each kernel. Each **cubin** file targets a specific compute-capability version and is forward-compatible *only with GPU architectures of the same major version number*. For example, **cubin** files that target compute capability 2.0 are supported on all compute-capability 2.x (Fermi) devices but are *not* supported on compute-capability 3.x (Kepler) devices. For this reason, to ensure forward compatibility with GPU architectures introduced after the application has been released, it is recommended that all applications support launching PTX versions of their kernels.¹

Applications that already include PTX versions of their kernels should work as-is on Kepler-based GPUs. Applications that only support specific GPU architectures via **cubin** files, however, will need to be updated to provide Kepler-compatible PTX or **cubins**.

¹ CUDA Runtime applications containing both cubin and PTX code for a given architecture will automatically use the cubin by default, keeping the PTX path strictly for forward-compatibility purposes.

1.3 Verifying Kepler Compatibility for Existing Applications

The first step is to check that Kepler-compatible device code (at least PTX) is compiled in to the application. The following sections show how to accomplish this for applications built with different CUDA Toolkit versions.

1.3.1 Applications Using CUDA Toolkit 4.1 or Earlier

CUDA applications built using CUDA Toolkit versions 2.1 through 4.1 are compatible with Kepler as long as they are built to include PTX versions of their kernels. To test that PTX JIT is working for your application, you can do the following:

- ▶ Download and install the latest driver from <http://www.nvidia.com/drivers>.
- ▶ Set the environment variable `CUDA_FORCE_PTX_JIT=1`.
- ▶ Create an empty temporary directory on your system.
- ▶ Set the environment variable `CUDA_CACHE_PATH` to be the path to this empty directory.
- ▶ Launch your application.

When starting a CUDA application for the first time with the above environment flag, the CUDA driver will JIT-compile the PTX for each CUDA kernel that is used into native **cubin** code. The generated **cubin** for the target GPU architecture is cached on disk by the CUDA driver.

If you set the environment variables above and then launch your program and it works properly, and if the directory you specified with the `CUDA_CACHE_PATH` environment variable is now populated with cache files, then you have successfully verified Kepler compatibility. Note that it is not necessary to inspect the contents of the cache files themselves; just check that the previously empty cache directory is now non-empty.

Be sure to unset these two environment variables when you are done testing if you do not normally use them. The temporary cache directory you created is safe to delete.

1.3.2 Applications Using CUDA Toolkit 5.0

CUDA applications built using CUDA Toolkit 5.0 are compatible with Kepler as long as they are built to include kernels in either Kepler-native **cubin** format (see [Building Applications with Kepler Support](#)) or PTX format (see [Applications Using CUDA Toolkit 4.1 or Earlier](#)) or both.

1.4 Building Applications with Kepler Support

When a CUDA application launches a kernel, the CUDA Runtime determines the compute capability of each GPU in the system and uses this information to automatically find the best matching **cubin** or PTX version of the kernel that is

available. If a **cubin** file supporting the architecture of the target GPU is available, it is used; otherwise, the CUDA Runtime will load the PTX and JIT-compile that PTX to the GPU's native **cubin** format before launching it. If neither is available, then the kernel launch will fail.

The method used to build your application with either native **cubin** or at least PTX support for Kepler depend on the version of the CUDA Toolkit used.

The main advantages of providing native **cubins** are as follows:

- ▶ It saves the end user the time it takes to PTX JIT a kernel that has been compiled as PTX. (However, since the CUDA driver will cache the **cubin** generated as a result of the PTX JIT, this is mostly a one-time cost for a given user.)
- ▶ PTX JIT-compiled kernels often cannot take advantage of architectural features of newer GPUs, meaning that native-compiled code may be faster or of greater accuracy.

1.4.1 Applications Using CUDA Toolkit 4.1 or Earlier

The compilers included in CUDA Toolkit 4.1 or earlier generate **cubin** files native to earlier NVIDIA architectures such as Fermi, but they *cannot* generate **cubin** files native to the Kepler architecture. To allow support for Kepler and future architectures when using version 4.1 or earlier of the CUDA Toolkit, the compiler must generate a PTX version of each kernel.

Below are compiler settings that could be used to build **mykernel.cu** to run on Fermi and earlier devices natively and on Kepler devices via PTX JIT.

Note that **compute_XX** refers to a PTX version and **sm_XX** refers to a **cubin** version. The **arch=** clause of the **-gencode=** command-line option to **nvcc** specifies the front-end compilation target and must always be a PTX version. The **code=** clause specifies the back-end compilation target and can either be **cubin** or PTX or both. **Only the back-end target version(s) specified by the code= clause will be retained in the resulting binary; at least one must be PTX to provide Kepler compatibility.**

Windows

```
nvcc.exe -ccbin "C:\vs2008\VC\bin"
-Xcompiler "/EHsc /W3 /nologo /O2 /Zi /MT"
-gencode=arch=compute_10,code=sm_10
-gencode=arch=compute_20,code=sm_20
-gencode=arch=compute_20,code=compute_20
--compile -o "Release\mykernel.cu.obj" "mykernel.cu"
```

Mac/Linux

```
/usr/local/cuda/bin/nvcc
-gencode=arch=compute_10,code=sm_10
-gencode=arch=compute_20,code=sm_20
-gencode=arch=compute_20,code=compute_20
-O2 -o mykernel.o -c mykernel.cu
```

Alternatively, you may be familiar with the simplified **nvcc** command-line option **-arch=sm_XX**, which is a shorthand equivalent to the following more explicit **-gencode=** command-line options used above. **-arch=sm_XX** expands to the following:

```
-gencode=arch=compute_XX,code=sm_XX
-gencode=arch=compute_XX,code=compute_XX
```

However, while the **-arch=sm_XX** command-line option does result in inclusion of a PTX back-end target by default, it can only specify a single target **cubin** architecture at a time, and it is not possible to use multiple **-arch=** options on the same **nvcc** command line, which is why the examples above use **-gencode=** explicitly.

1.4.2 Applications Using CUDA Toolkit 5.0

With version 5.0 of the CUDA Toolkit, **nvcc** can generate **cubin** files native to the Kepler architecture (compute capability 3.x). When using CUDA Toolkit 5.0, to ensure that **nvcc** will generate **cubin** files for all released GPU architectures as well as a PTX version for forward compatibility with future GPU architectures, specify the appropriate **-gencode=** parameters on the **nvcc** command line as shown in the examples below.

Windows

```
nvcc.exe -ccbin "C:\vs2008\VC\bin"
-Xcompiler "/EHsc /W3 /nologo /O2 /Zi /MT"
-gencode=arch=compute_10,code=sm_10
-gencode=arch=compute_20,code=sm_20
-gencode=arch=compute_30,code=sm_30
-gencode=arch=compute_35,code=sm_35
-gencode=arch=compute_35,code=compute_35
--compile -o "Release\mykernel.cu.obj" "mykernel.cu"
```

Mac/Linux

```
/usr/local/cuda/bin/nvcc
-gencode=arch=compute_10,code=sm_10
-gencode=arch=compute_20,code=sm_20
-gencode=arch=compute_30,code=sm_30
-gencode=arch=compute_35,code=sm_35
-gencode=arch=compute_35,code=compute_35
-O2 -o mykernel.o -c mykernel.cu
```

Note that **compute_XX** refers to a PTX version and **sm_XX** refers to a **cubin** version. The **arch=** clause of the **-gencode=** command-line option to **nvcc** specifies the front-end compilation target and must always be a PTX version. The **code=** clause specifies the back-end compilation target and can either be **cubin** or PTX or both. **Only the back-end target version(s) specified by the code= clause will be retained in the resulting binary; at least one should be PTX to provide compatibility with future architectures.**

Appendix A.

REVISION HISTORY

Version 1.0

- ▶ Initial public release.

Version 1.1

- ▶ Added `sm_35`.
- ▶ Simplified by removing discussion of CUDA Driver API, which is infrequently used

Notice

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.

Information furnished is believed to be accurate and reliable. However, NVIDIA Corporation assumes no responsibility for the consequences of use of such information or for any infringement of patents or other rights of third parties that may result from its use. No license is granted by implication of otherwise under any patent rights of NVIDIA Corporation. Specifications mentioned in this publication are subject to change without notice. This publication supersedes and replaces all other information previously supplied. NVIDIA Corporation products are not authorized as critical components in life support devices or systems without express written approval of NVIDIA Corporation.

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

© 2007-2012 NVIDIA Corporation. All rights reserved.