



NVIDIA CUDA™

Programming Guide

Version 2.3.1

8/26/2009

Table of Contents

Chapter 1. Introduction	1
1.1 From Graphics Processing to General-Purpose Parallel Computing	1
1.2 CUDA™: a General-Purpose Parallel Computing Architecture	3
1.3 CUDA's Scalable Programming Model.....	4
1.4 Document's Structure	5
Chapter 2. Programming Model.....	7
2.1 Kernels	7
2.2 Thread Hierarchy.....	8
2.3 Memory Hierarchy	10
2.4 Host and Device	11
2.5 Compute Capability	14
Chapter 3. Programming Interface	15
3.1 Compilation with NVCC	16
3.1.1 <code>__noinline__</code>	16
3.1.2 <code>#pragma unroll</code>	17
3.2 C for CUDA	17
3.2.1 Device Memory.....	17
3.2.2 Shared Memory	20
3.2.3 Multiple Devices.....	26
3.2.4 Texture Memory	27
3.2.4.1 Texture Reference Declaration	28
3.2.4.2 Runtime Texture Reference Attributes	28
3.2.4.3 Texture Binding	29
3.2.5 Page-Locked Host Memory	32
3.2.5.1 Portable Memory.....	32
3.2.5.2 Write-Combining Memory	32
3.2.5.3 Mapped Memory	32
3.2.6 Asynchronous Concurrent Execution	33

3.2.6.1	Stream	34
3.2.6.2	Event	35
3.2.6.3	Synchronous Calls	35
3.2.7	OpenGL Interoperability	36
3.2.8	Direct3D Interoperability	38
3.2.9	Error Handling	43
3.2.10	Debugging using the Device Emulation Mode	44
3.3	Driver API	45
3.3.1	Context	47
3.3.2	Module	48
3.3.3	Kernel Execution	49
3.3.4	Device Memory	50
3.3.5	Shared Memory	53
3.3.6	Multiple Devices	54
3.3.7	Texture Memory	55
3.3.8	Page-Locked Host Memory	57
3.3.9	Asynchronous Concurrent Execution	57
3.3.9.1	Stream	57
3.3.9.2	Event Management	58
3.3.9.3	Synchronous Calls	59
3.3.10	OpenGL Interoperability	59
3.3.11	Direct3D Interoperability	61
3.3.12	Error Handling	67
3.4	Versioning and Compatibility	67
3.5	Compute Modes	68
3.6	Mode Switches	69
Chapter 4. Hardware Implementation		71
4.1	A Set of SIMT Multiprocessors with On-Chip Shared Memory	71
4.2	Multiple Devices	75
Chapter 5. Performance Guidelines		77
5.1	Instruction Performance	77
5.1.1	Instruction Throughput	77
5.1.1.1	Arithmetic Instructions	77

5.1.1.2	Control Flow Instructions	79
5.1.1.3	Memory Instructions	80
5.1.1.4	Synchronization Instruction	80
5.1.2	Memory Bandwidth	80
5.1.2.1	Global Memory	81
5.1.2.2	Local Memory	88
5.1.2.3	Constant Memory	89
5.1.2.4	Texture Memory	89
5.1.2.5	Shared Memory	89
5.1.2.6	Registers	97
5.2	Execution Configuration	97
5.3	Data Transfer between Host and Device	98
5.4	Warp-Level Synchronization	99
5.5	Overall Performance Optimization Strategies	99
Appendix A. Technical Specifications		101
A.1	General Specifications	101
A.1.1	Specifications for Compute Capability 1.0	102
A.1.2	Specifications for Compute Capability 1.1	103
A.1.3	Specifications for Compute Capability 1.2	103
A.1.4	Specifications for Compute Capability 1.3	103
A.2	Floating-Point Standard	103
Appendix B. C Extensions		105
B.1	Function Type Qualifiers	105
B.1.1	__device__	105
B.1.2	__global__	105
B.1.3	__host__	105
B.1.4	Restrictions	106
B.2	Variable Type Qualifiers	106
B.2.1	__device__	106
B.2.2	__constant__	106
B.2.3	__shared__	107
B.2.4	Volatile	107
B.2.5	Restrictions	108

B.3	Built-in Vector Types	109
B.3.1	char1, uchar1, char2, uchar2, char3, uchar3, char4, uchar4, short1, ushort1, short2, ushort2, short3, ushort3, short4, ushort4, int1, uint1, int2, uint2, int3, uint3, int4, uint4, long1, ulong1, long2, ulong2, long3, ulong3, long4, ulong4, longlong1, longlong2, float1, float2, float3, float4, double1, double2	109
B.3.2	dim3	110
B.4	Built-in Variables	110
B.4.1	gridDim	110
B.4.2	blockIdx	110
B.4.3	blockDim	110
B.4.4	threadIdx	110
B.4.5	warpSize	110
B.4.6	Restrictions	110
B.5	Memory Fence Functions	111
B.6	Synchronization Function	112
B.7	Mathematical Functions	112
B.8	Texture Functions	113
B.8.1	tex1Dfetch()	113
B.8.2	tex1D()	114
B.8.3	tex2D()	114
B.8.4	tex3D()	114
B.9	Time Function	114
B.10	Atomic Functions	115
B.10.1	Arithmetic Functions	115
B.10.1.1	atomicAdd()	115
B.10.1.2	atomicSub()	115
B.10.1.3	atomicExch()	115
B.10.1.4	atomicMin()	116
B.10.1.5	atomicMax()	116
B.10.1.6	atomicInc()	116
B.10.1.7	atomicDec()	116
B.10.1.8	atomicCAS()	116
B.10.2	Bitwise Functions	117
B.10.2.1	atomicAnd()	117

B.10.2.2	atomicOr().....	117
B.10.2.3	atomicXor()	117
B.11	Warp Vote Functions	117
B.12	Execution Configuration	118
Appendix C. Mathematical Functions		119
C.1	Standard Functions.....	119
C.1.1	Single-Precision Floating-Point Functions.....	119
C.1.2	Double-Precision Floating-Point Functions	121
C.1.3	Integer Functions	123
C.2	Intrinsic Functions	124
C.2.1	Single-Precision Floating-Point Functions.....	124
C.2.2	Double-Precision Floating-Point Functions	125
C.2.3	Integer Functions	126
Appendix D. C++ Language Constructs		129
D.1	Polymorphism	129
D.2	Default Parameters.....	130
D.3	Operator Overloading	130
D.4	Namespaces.....	131
D.5	Function Templates	131
Appendix E. Texture Fetching.....		133
E.1	Nearest-Point Sampling.....	134
E.2	Linear Filtering	134
E.3	Table Lookup	136

List of Figures

Figure 1-1. Floating-Point Operations per Second and Memory Bandwidth for the CPU and GPU 2	
Figure 1-2. The GPU Devotes More Transistors to Data Processing	3
Figure 1-3. CUDA is Designed to Support Various Languages or Application Programming Interfaces	4
Figure 2-1. Grid of Thread Blocks.....	10
Figure 2-2. Memory Hierarchy	11
Figure 2-3. Heterogeneous Programming	13
Figure 3-1. Matrix Multipliation without Shared Memory	22
Figure 3-2. Matrix Multipliation with Shared Memory	26
Figure 3-3. Library Context Management.....	48
Figure 3-4. The Driver API is Forward Compatible	68
Figure 4-1. Automatic Scalability	72
Figure 4-2. Hardware Model	75
Figure 5-1. Examples of Coalesced Global Memory Access Patterns.....	84
Figure 5-2. Examples of Global Memory Access Patterns That Are Non-Coalesced for Devices of Compute Capability 1.0 or 1.1	85
Figure 5-3. Examples of Global Memory Access Patterns That Are Non-Coalesced for Devices of Compute Capability 1.0 or 1.1	86
Figure 5-4. Examples of Global Memory Access by Devices with Compute Capability 1.2 and Higher	87
Figure 5-5. Examples of Shared Memory Access Patterns without Bank Conflicts	93
Figure 5-6. Example of a Shared Memory Access Pattern without Bank Conflicts.....	94
Figure 5-7. Examples of Shared Memory Access Patterns with Bank Conflicts	95
Figure 5-8. Example of Shared Memory Read Access Patterns with Broadcast.....	96



Chapter 1. Introduction

1.1 From Graphics Processing to General-Purpose Parallel Computing

Driven by the insatiable market demand for realtime, high-definition 3D graphics, the programmable Graphic Processor Unit or GPU has evolved into a highly parallel, multithreaded, manycore processor with tremendous computational horsepower and very high memory bandwidth, as illustrated by Figure 1-1.

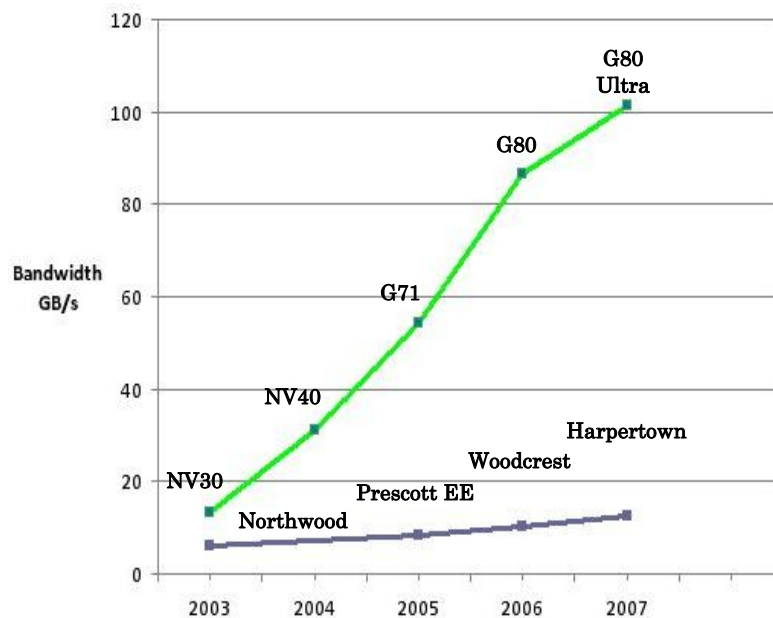
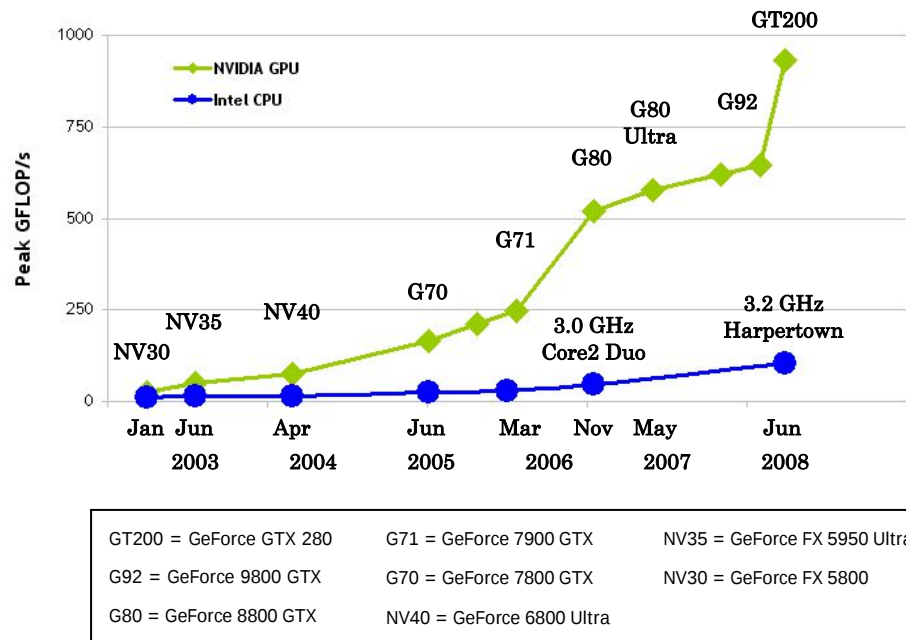


Figure 1-1. Floating-Point Operations per Second and Memory Bandwidth for the CPU and GPU

The reason behind the discrepancy in floating-point capability between the CPU and the GPU is that the GPU is specialized for compute-intensive, highly parallel computation – exactly what graphics rendering is about – and therefore designed such that more transistors are devoted to data processing rather than data caching and flow control, as schematically illustrated by Figure 1-2.

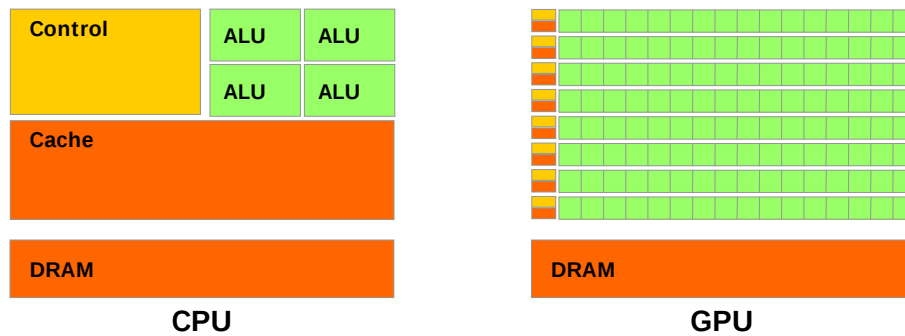


Figure 1-2. The GPU Devotes More Transistors to Data Processing

More specifically, the GPU is especially well-suited to address problems that can be expressed as data-parallel computations – the same program is executed on many data elements in parallel – with high arithmetic intensity – the ratio of arithmetic operations to memory operations. Because the same program is executed for each data element, there is a lower requirement for sophisticated flow control; and because it is executed on many data elements and has high arithmetic intensity, the memory access latency can be hidden with calculations instead of big data caches.

Data-parallel processing maps data elements to parallel processing threads. Many applications that process large data sets can use a data-parallel programming model to speed up the computations. In 3D rendering, large sets of pixels and vertices are mapped to parallel threads. Similarly, image and media processing applications such as post-processing of rendered images, video encoding and decoding, image scaling, stereo vision, and pattern recognition can map image blocks and pixels to parallel processing threads. In fact, many algorithms outside the field of image rendering and processing are accelerated by data-parallel processing, from general signal processing or physics simulation to computational finance or computational biology.

1.2 CUDA™: a General-Purpose Parallel Computing Architecture

In November 2006, NVIDIA introduced CUDA™, a general purpose parallel computing architecture – with a new parallel programming model and instruction set architecture – that leverages the parallel compute engine in NVIDIA GPUs to solve many complex computational problems in a more efficient way than on a CPU.

CUDA comes with a software environment that allows developers to use C as a high-level programming language. As illustrated by Figure 1-3, other languages or application programming interfaces will be supported in the future, such as FORTRAN, C++, OpenCL, and DirectX Compute.

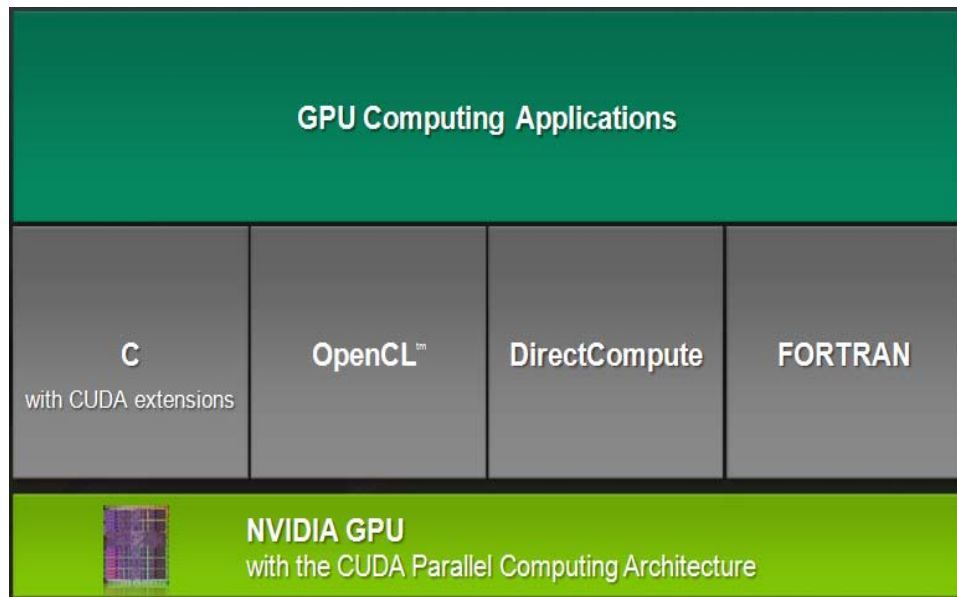


Figure 1-3. CUDA is Designed to Support Various Languages or Application Programming Interfaces

1.3 CUDA's Scalable Programming Model

The advent of multicore CPUs and manycore GPUs means that mainstream processor chips are now parallel systems. Furthermore, their parallelism continues to scale with Moore's law. The challenge is to develop application software that transparently scales its parallelism to leverage the increasing number of processor cores, much as 3D graphics applications transparently scale their parallelism to manycore GPUs with widely varying numbers of cores.

CUDA's parallel programming model is designed to overcome this challenge while maintaining a low learning curve for programmers familiar with standard programming languages such as C.

At its core are three key abstractions – a hierarchy of thread groups, shared memories, and barrier synchronization – that are simply exposed to the programmer as a minimal set of language extensions.

These abstractions provide fine-grained data parallelism and thread parallelism, nested within coarse-grained data parallelism and task parallelism. They guide the programmer to partition the problem into coarse sub-problems that can be solved independently in parallel, and then into finer pieces that can be solved cooperatively in parallel. Such a decomposition preserves language expressivity by allowing threads to cooperate when solving each sub-problem, and at the same time enables transparent scalability since each sub-problem can be scheduled to be solved on any of the available processor cores: A compiled CUDA program can therefore execute on any number of processor cores, and only the runtime system needs to know the physical processor count.

This scalable programming model allows the CUDA architecture to span a wide market range by simply scaling the number of processors and memory partitions: from the high-performance enthusiast GeForce GTX 280 GPU and professional Quadro and Tesla computing products to a variety of inexpensive, mainstream GeForce GPUs (see Appendix A for a list of all CUDA-enabled GPUs).

1.4 Document's Structure

This document is organized into the following chapters:

- ❑ Chapter 1 is a general introduction to CUDA.
- ❑ Chapter 2 outlines CUDA's programming model.
- ❑ Chapter 3 describes the programming interface.
- ❑ Chapter 4 describes the hardware implementation.
- ❑ Chapter 5 gives some guidance on how to achieve maximum performance.
- ❑ Appendix A lists the technical specifications of various devices.
- ❑ Appendix B is a detailed description of all extensions to the C language.
- ❑ Appendix C lists the mathematical functions supported in CUDA.
- ❑ Appendix D lists the C++ constructs supported in device code.
- ❑ Appendix E gives more details on texture fetching.

Chapter 2.

Programming Model

This chapter introduces the main concepts that make up the CUDA programming model by outlining how they are exposed in C. An extensive description of C for CUDA is given in Section 3.2.

2.1 Kernels

C for CUDA extends C by allowing the programmer to define C functions, called *kernels*, that, when called, are executed N times in parallel by N different *CUDA threads*, as opposed to only once like regular C functions.

A kernel is defined using the `__global__` declaration specifier and the number of CUDA threads for each call is specified using a new `<<<...>>>` syntax:

```
// Kernel definition
__global__ void VecAdd(float* A, float* B, float* C)
{
    ...
}

int main()
{
    ...
    // Kernel invocation
    VecAdd<<<1, N>>>(A, B, C);
}
```

Each of the threads that execute a kernel is given a unique *thread ID* that is accessible within the kernel through the built-in `threadIdx` variable. As an illustration, the following sample code adds two vectors A and B of size N and stores the result into vector C :

```
// Kernel definition
__global__ void VecAdd(float* A, float* B, float* C)
{
    int i = threadIdx.x;
    C[i] = A[i] + B[i];
}

int main()
{
```

```

...
// Kernel invocation
VecAdd<<<1, N>>>(A, B, C);
}

```

Each of the threads that execute **VecAdd()** performs one pair-wise addition.

2.2 Thread Hierarchy

For convenience, **threadIdx** is a 3-component vector, so that threads can be identified using a one-dimensional, two-dimensional, or three-dimensional *thread index*, forming a one-dimensional, two-dimensional, or three-dimensional *thread block*. This provides a natural way to invoke computation across the elements in a domain such as a vector, matrix, or field. As an example, the following code adds two matrices *A* and *B* of size $N \times N$ and stores the result into matrix *C*:

```

// Kernel definition
__global__ void MatAdd(float A[N][N], float B[N][N],
                      float C[N][N])
{
    int i = threadIdx.x;
    int j = threadIdx.y;
    C[i][j] = A[i][j] + B[i][j];
}

int main()
{
    ...
    // Kernel invocation
    dim3 dimBlock(N, N);
    MatAdd<<<1, dimBlock>>>(A, B, C);
}

```

The index of a thread and its thread ID relate to each other in a straightforward way: For a one-dimensional block, they are the same; for a two-dimensional block of size (D_x, D_y) , the thread ID of a thread of index (x, y) is $(x + y D_x)$; for a three-dimensional block of size (D_x, D_y, D_z) , the thread ID of a thread of index (x, y, z) is $(x + y D_x + z D_x D_y)$.

Threads within a block can cooperate among themselves by sharing data through some *shared memory* and synchronizing their execution to coordinate memory accesses. More precisely, one can specify synchronization points in the kernel by calling the **__syncthreads()** intrinsic function; **__syncthreads()** acts as a barrier at which all threads in the block must wait before any is allowed to proceed. Section 3.2.2 gives an example of using shared memory.

For efficient cooperation, the shared memory is expected to be a low-latency memory near each processor core, much like an L1 cache, **__syncthreads()** is expected to be lightweight, and all threads of a block are expected to reside on the same processor core. The number of threads per block is therefore restricted by the limited memory resources of a processor core. On current GPUs, a thread block may contain up to 512 threads.

However, a kernel can be executed by multiple equally-shaped thread blocks, so that the total number of threads is equal to the number of threads per block times the

number of blocks. These multiple blocks are organized into a one-dimensional or two-dimensional *grid* of thread blocks as illustrated by Figure 2-1. The dimension of the grid is specified by the first parameter of the `<<<...>>>` syntax. Each block within the grid can be identified by a one-dimensional or two-dimensional index accessible within the kernel through the built-in **blockIdx** variable. The dimension of the thread block is accessible within the kernel through the built-in **blockDim** variable. The previous sample code becomes:

```
// Kernel definition
__global__ void MatAdd(float A[N][N], float B[N][N],
                      float C[N][N])
{
    int i = blockIdx.x * blockDim.x + threadIdx.x;
    int j = blockIdx.y * blockDim.y + threadIdx.y;
    if (i < N && j < N)
        C[i][j] = A[i][j] + B[i][j];
}

int main()
{
    ...
    // Kernel invocation
    dim3 dimBlock(16, 16);
    dim3 dimGrid((N + dimBlock.x - 1) / dimBlock.x,
                 (N + dimBlock.y - 1) / dimBlock.y);
    MatAdd<<<dimGrid, dimBlock>>>(A, B, C);
}
```

The thread block size of $16 \times 16 = 256$ threads was chosen somewhat arbitrarily, and a grid is created with enough blocks to have one thread per matrix element as before.

Thread blocks are required to execute independently: It must be possible to execute them in any order, in parallel or in series. This independence requirement allows thread blocks to be scheduled in any order across any number of cores, enabling programmers to write code that scales with the number of cores.

The number of thread blocks in a grid is typically dictated by the size of the data being processed rather than by the number of processors in the system, which it can greatly exceed.

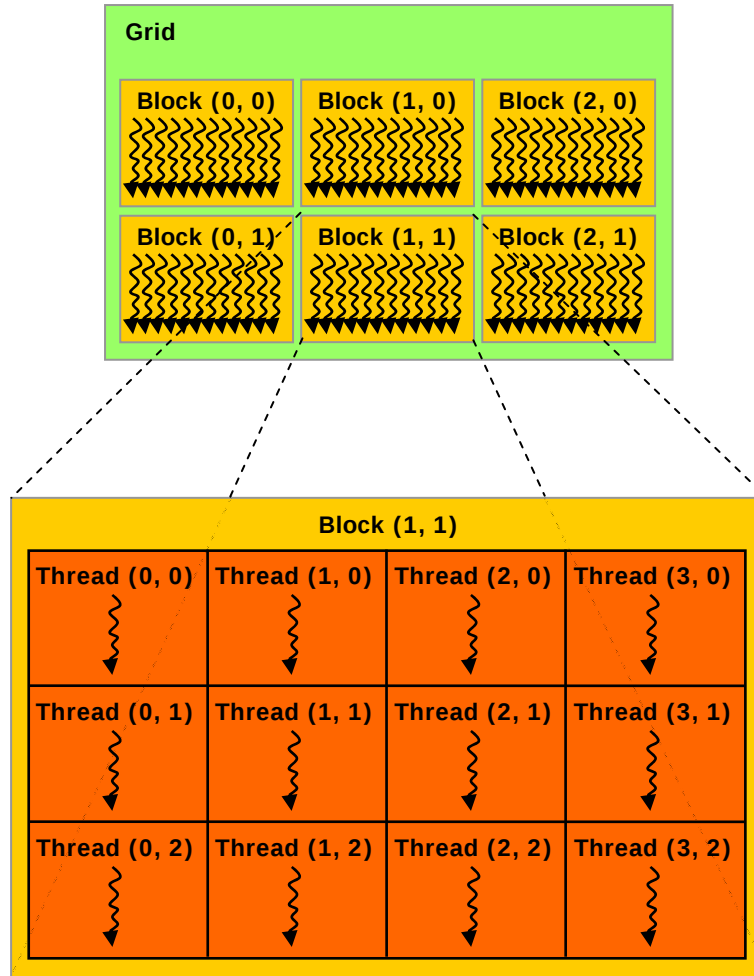


Figure 2-1. Grid of Thread Blocks

2.3 Memory Hierarchy

CUDA threads may access data from multiple memory spaces during their execution as illustrated by Figure 2-2. Each thread has a private local memory. Each thread block has a shared memory visible to all threads of the block and with the same lifetime as the block. Finally, all threads have access to the same global memory.

There are also two additional read-only memory spaces accessible by all threads: the constant and texture memory spaces. The global, constant, and texture memory spaces are optimized for different memory usages (see Sections 5.1.2.1, 5.1.2.3, and 5.1.2.4). Texture memory also offers different addressing modes, as well as data filtering, for some specific data formats (see Section 3.2.4).

The global, constant, and texture memory spaces are persistent across kernel launches by the same application.

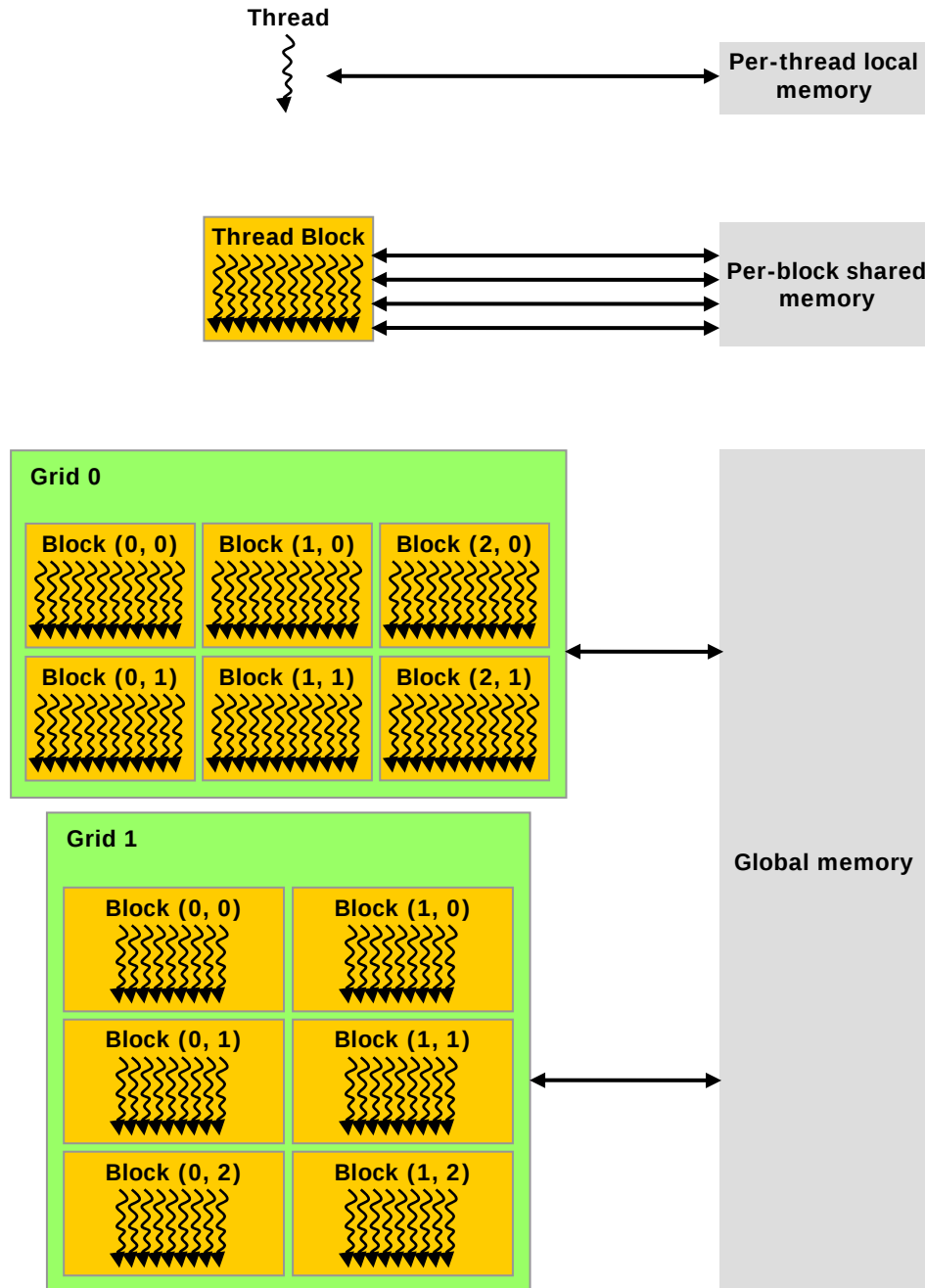
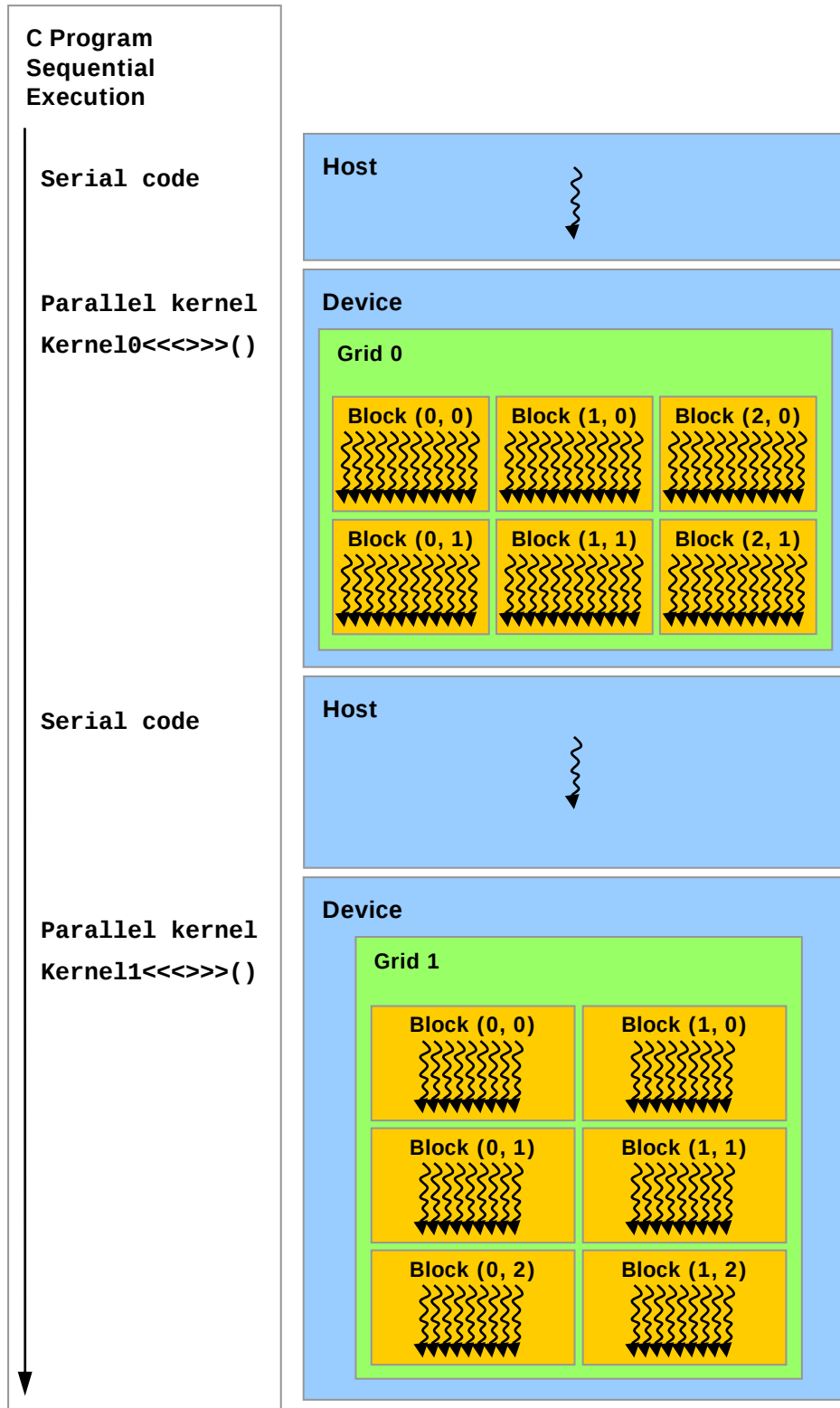


Figure 2-2. Memory Hierarchy

2.4 Host and Device

As illustrated by Figure 2-3, CUDA's programming model assumes that the CUDA threads execute on a physically separate *device* that operates as a coprocessor to the *host* running the C program. This is the case, for example, when the kernels execute on a GPU and the rest of the C program executes on a CPU.

CUDA's programming model also assumes that both the host and the device maintain their own DRAM, referred to as *host memory* and *device memory*, respectively. Therefore, a program manages the global, constant, and texture memory spaces visible to kernels through calls to the CUDA runtime (described in Chapter 3). This includes device memory allocation and deallocation, as well as data transfer between host and device memory.



Serial code executes on the host while parallel code executes on the device.

Figure 2-3. Heterogeneous Programming

2.5 Compute Capability

The *compute capability* of a device is defined by a major revision number and a minor revision number.

Devices with the same major revision number are of the same core architecture. The devices listed in Appendix A are all of compute capability 1.x (Their major revision number is 1).

The minor revision number corresponds to an incremental improvement to the core architecture, possibly including new features.

The technical specifications of the various compute capabilities are given in Appendix A.

Chapter 3. Programming Interface

Two interfaces are currently supported to write CUDA programs: *C for CUDA* and the *CUDA driver API*. They are mutually exclusive: A program must use either one or the other.

C for CUDA exposes the CUDA programming model as a minimal set of extensions to the C language. Any source file that contains some of these extensions must be compiled with **nvcc** as outlined in Section 3.1. These extensions allow programmers to define a kernel as a C function and use some new syntax to specify the grid and block dimension each time the function is called.

The CUDA driver API is a lower-level C API that provides functions to load kernels as modules of CUDA binary or assembly code, to inspect their parameters, and to launch them. Binary or assembly code are usually obtained by compiling kernels written in C.

C for CUDA comes with a *runtime API* and both the runtime API and the driver API provide functions to allocate and deallocate device memory, transfer data between host memory and device memory, manage systems with multiple devices, etc.

The runtime API is built on top of the CUDA driver API. Initialization, context, and module management are all implicit and resulting code is more concise. C for CUDA also supports device emulation, which facilitates debugging (see Section 3.2.9).

In contrast, the CUDA driver API requires more code, is harder to program and debug, but offers a better level of control and is language-independent since it handles binary or assembly code.

Section 3.2 continues the description of C for CUDA started in Chapter 2. It also introduces concepts that are common to both C for CUDA and the driver API: linear memory, CUDA arrays, shared memory, texture memory, page-locked host memory, device enumeration, asynchronous execution, interoperability with graphics APIs. Section 3.3 assumes knowledge of these concepts and describes how they are exposed by the driver API.

3.1 Compilation with NVCC

Kernels can be written using CUDA's instruction set architecture, called *PTX*, which is described in a separate document. It is however usually more effective to use a high-level programming language such as C. In both cases, kernels must be compiled into binary code by **nvcc**.

nvcc is a compiler driver that simplifies the process of compiling C for CUDA code: It provides simple and familiar command line options and executes them by invoking the collection of tools that implement the different compilation stages.

Source files can include a mix of host code (i.e. code that executes on the host) and device code (i.e. code that executes on the device). **nvcc**'s basic workflow consists in separating device code from host code and compiling the device code into an assembly form (*PTX* code) or binary form (*cubin* object). The generated host code is output either as C code that is left to be compiled using another tool or as object code directly by invoking the host compiler during the last compilation stage.

Applications can either ignore the generated host code (if any) and load and execute the *PTX* code or *cubin* object on the device using the CUDA driver API (see Section 3.3), or they can link to the generated host code, which includes the *cubin* object as a global initialized data array and contains a translation of the execution configuration syntax described in Section B.12 into the necessary C for CUDA runtime startup code to load and launch each compiled kernel.

The front end of the compiler processes CUDA source files according to C++ syntax rules. Full C++ is supported for the host code. However, only the C subset of C++ is fully supported for the device code; C++ specific features such as classes, inheritance, or declaration of variables within basic blocks are not. As a consequence of the use of C++ syntax rules, void pointers (e.g. returned by **malloc()**) cannot be assigned to non-void pointers without a typecast.

nvcc introduces two compiler directives described in the following sections.

Some *PTX* instructions are only supported on devices of higher compute capabilities. For example, atomic instructions on global memory are only supported on devices of compute capability 1.1 and above; double-precision instructions are only supported on devices of compute capability 1.3 and above. The **-arch** compiler option specifies the compute capability that is assumed when compiling to *PTX* code. So, code that contains double-precision arithmetic, for example, must be compiled with "**-arch sm_13**" (or higher compute capability), otherwise double-precision arithmetic will get demoted to single-precision arithmetic.

A detailed description of **nvcc**'s workflow and command options can be found in a separate document.

3.1.1 `__noinline__`

By default, a `__device__` function is always inlined. The `__noinline__` function qualifier however can be used as a hint for the compiler not to inline the function if possible. The function body must still be in the same file where it is called.

The compiler will not honor the `__noinline__` qualifier for functions with pointer parameters and for functions with large parameter lists.

3.1.2 `#pragma unroll`

By default, the compiler unrolls small loops with a known trip count. The **`#pragma unroll`** directive however can be used to control unrolling of any given loop. It must be placed immediately before the loop and only applies to that loop. It is optionally followed by a number that specifies how many times the loop must be unrolled.

For example, in this code sample:

```
#pragma unroll 5
for (int i = 0; i < n; ++i)
```

the loop will be unrolled 5 times. The compiler will also insert code to ensure correctness (in the example above, to ensure that there will only be **`n`** iterations if **`n`** is less than 5, for example). It is up to the programmer to make sure that the specified unroll number gives the best performance.

`#pragma unroll 1` will prevent the compiler from ever unrolling a loop.

If no number is specified after **`#pragma unroll`**, the loop is completely unrolled if its trip count is constant, otherwise it is not unrolled at all.

3.2 C for CUDA

C for CUDA provides a simple path for users familiar with the C programming language to easily write programs for execution by the device.

It consists of a minimal set of extensions to the C language and a runtime library. The core language extensions have been introduced in Chapter 2. This section continues with an introduction to the runtime. A complete description of all extensions can be found in Appendix B and a complete description of the runtime in the CUDA reference manual.

The runtime is implemented in the **`cuda`** dynamic library and all its entry points are prefixed with **`cuda`**.

There is no explicit initialization function for the runtime; it initializes the first time a runtime function is called. One needs to keep this in mind when timing runtime function calls and when interpreting the error code from the first call into the runtime.

On system with multiple devices, kernels are executed on device 0 by default as detailed in Section 3.2.3.

3.2.1 Device Memory

As mentioned in Section 2.4, CUDA's programming model assumes a system composed of a host and a device, each with their own separate memory. Kernels can only operate out of device memory, so the runtime provides functions to

allocate, deallocate, and copy device memory, as well as transfer data between host memory and device memory.

Device memory can be allocated either as *linear memory* or as *CUDA arrays*.

CUDA arrays are opaque memory layouts optimized for texture fetching. They are described in Section 3.2.4.

Linear memory exists on the device in a 32-bit address space, so separately allocated entities can reference one another via pointers, for example, in a binary tree.

Linear memory is typically allocated using **cudaMalloc()** and freed using **cudaFree()** and data transfer between host memory and device memory are typically done using **cudaMemcpy()**. In the vector addition code sample of Section 2.1, the vectors need to be copied from host memory to device memory:

```
// Device code
__global__ void VecAdd(float* A, float* B, float* C)
{
    int i = blockDim.x * blockIdx.x + threadIdx.x;
    if (i < N)
        C[i] = A[i] + B[i];
}

// Host code
int main()
{
    int N = ...;
    size_t size = N * sizeof(float);

    // Allocate input vectors h_A and h_B in host memory
    float* h_A = malloc(size);
    float* h_B = malloc(size);

    // Allocate vectors in device memory
    float* d_A;
    cudaMalloc((void**)&d_A, size);
    float* d_B;
    cudaMalloc((void**)&d_B, size);
    float* d_C;
    cudaMalloc((void**)&d_C, size);

    // Copy vectors from host memory to device memory
    cudaMemcpy(d_A, h_A, size, cudaMemcpyHostToDevice);
    cudaMemcpy(d_B, h_B, size, cudaMemcpyHostToDevice);

    // Invoke kernel
    int threadsPerBlock = 256;
    int blocksPerGrid =
        (N + threadsPerBlock - 1) / threadsPerBlock;
    VecAdd<<<blocksPerGrid, threadsPerBlock>>>>(d_A, d_B, d_C);

    // Copy result from device memory to host memory
    // h_C contains the result in host memory
    cudaMemcpy(h_C, d_C, size, cudaMemcpyDeviceToHost);

    // Free device memory
    cudaFree(d_A);
```

```

    cudaFree(d_B);
    cudaFree(d_C);
}

```

Linear memory can also be allocated through **cudaMallocPitch()** and **cudaMalloc3D()**. These functions are recommended for allocations of 2D or 3D arrays as it makes sure that the allocation is appropriately padded to meet the alignment requirements described in Section 5.1.2.1, therefore ensuring best performance when accessing the row addresses or performing copies between 2D arrays and other regions of device memory (using the **cudaMemcpy2D()** and **cudaMemcpy3D()** functions). The returned pitch (or stride) must be used to access array elements. The following code sample allocates a **width×height** 2D array of floating-point values and shows how to loop over the array elements in device code:

```

// Host code
float* devPtr;
int pitch;
cudaMallocPitch((void**)&devPtr, &pitch,
               width * sizeof(float), height);
myKernel<<<100, 512>>>(devPtr, pitch);

// Device code
__global__ void myKernel(float* devPtr, int pitch)
{
    for (int r = 0; r < height; ++r) {
        float* row = (float*)((char*)devPtr + r * pitch);
        for (int c = 0; c < width; ++c) {
            float element = row[c];
        }
    }
}

```

The following code sample allocates a **width×height×depth** 3D array of floating-point values and shows how to loop over the array elements in device code:

```

// Host code
cudaPitchedPtr devPitchedPtr;
cudaExtent extent = make_cudaExtent(64, 64, 64);
cudaMalloc3D(&devPitchedPtr, extent);
myKernel<<<100, 512>>>(devPitchedPtr, extent);

// Device code
__global__ void myKernel(cudaPitchedPtr devPitchedPtr,
                        cudaExtent extent)
{
    char* devPtr = devPitchedPtr.ptr;
    size_t pitch = devPitchedPtr.pitch;
    size_t slicePitch = pitch * extent.height;
    for (int z = 0; z < extent.depth; ++z) {
        char* slice = devPtr + z * slicePitch;
        for (int y = 0; y < extent.height; ++y) {
            float* row = (float*)(slice + y * pitch);
            for (int x = 0; x < extent.width; ++x) {
                float element = row[x];
            }
        }
    }
}

```

The reference manual lists all the various functions used to copy memory between linear memory allocated with **cudaMalloc()**, linear memory allocated with **cudaMallocPitch()** or **cudaMalloc3D()**, CUDA arrays, and memory allocated for variables declared in global or constant memory space.

The following code sample copies the 2D array to the CUDA array allocated in the previous code samples:

```
cudaMemcpy2DToArray(cuArray, 0, 0, devPtr, pitch,
                    width * sizeof(float), height,
                    cudaMemcpyDeviceToDevice);
```

The following code sample copies some host memory array to constant memory:

```
__constant__ float constData[256];
float data[256];
cudaMemcpyToSymbol(constData, data, sizeof(data));
```

cudaGetSymbolAddress() is used to retrieve the address pointing to the memory allocated for a variable declared in global memory space. The size of the allocated memory is obtained through **cudaGetSymbolSize()**.

3.2.2 Shared Memory

As detailed in Section B.2 shared memory is allocated using the **__shared__** qualifier.

Shared memory is expected to be much faster than global memory as mentioned in Section 2.2 and detailed in Section 5.1.2.5. Any opportunity to replace global memory accesses by shared memory accesses should therefore be exploited as illustrated by the following matrix multiplication example.

The following code sample is a straightforward implementation of matrix multiplication that does not take advantage of shared memory. Each thread reads one row of *A* and one column of *B* and computes the corresponding element of *C* as illustrated in Figure 3-1. *A* is therefore read *B.width* times from global memory and *B* is read *A.height* times.

```
// Matrices are stored in row-major order:
// M(row, col) = *(M.elements + row * M.width + col)
typedef struct {
    int width;
    int height;
    float* elements;
} Matrix;

// Thread block size
#define BLOCK_SIZE 16

// Forward declaration of the matrix multiplication kernel
__global__ void MatMulKernel(const Matrix, const Matrix, Matrix);

// Matrix multiplication - Host code
// Matrix dimensions are assumed to be multiples of BLOCK_SIZE
void MatMul(const Matrix A, const Matrix B, Matrix C)
{
    // Load A and B to device memory
    Matrix d_A;
```

```

d_A.width = A.width; d_A.height = A.height;
size_t size = A.width * A.height * sizeof(float);
cudaMalloc((void**)&d_A.elements, size);
cudaMemcpy(d_A.elements, A.elements, size,
           cudaMemcpyHostToDevice);
Matrix d_B;
d_B.width = B.width; d_B.height = B.height;
size = B.width * B.height * sizeof(float);
cudaMalloc((void**)&d_B.elements, size);
cudaMemcpy(d_B.elements, B.elements, size,
           cudaMemcpyHostToDevice);

// Allocate C in device memory
Matrix d_C;
d_C.width = C.width; d_C.height = C.height;
size = C.width * C.height * sizeof(float);
cudaMalloc((void**)&d_C.elements, size);

// Invoke kernel
dim3 dimBlock(BLOCK_SIZE, BLOCK_SIZE);
dim3 dimGrid(B.width / dimBlock.x, A.height / dimBlock.y);
MatMulKernel<<<dimGrid, dimBlock>>>(d_A, d_B, d_C);

// Read C from device memory
cudaMemcpy(C.elements, Cd.elements, size,
           cudaMemcpyDeviceToHost);

// Free device memory
cudaFree(d_A.elements);
cudaFree(d_B.elements);
cudaFree(d_C.elements);
}

// Matrix multiplication kernel called by MatMul()
__global__ void MatMulKernel(Matrix A, Matrix B, Matrix C)
{
    // Each thread computes one element of C
    // by accumulating results into Cvalue
    float Cvalue = 0;
    int row = blockIdx.y * blockDim.y + threadIdx.y;
    int col = blockIdx.x * blockDim.x + threadIdx.x;
    for (int e = 0; e < A.width; ++e)
        Cvalue += A.elements[row * A.width + e]
                 * B.elements[e * B.width + col];
    C.elements[row * C.width + col] = Cvalue;
}

```

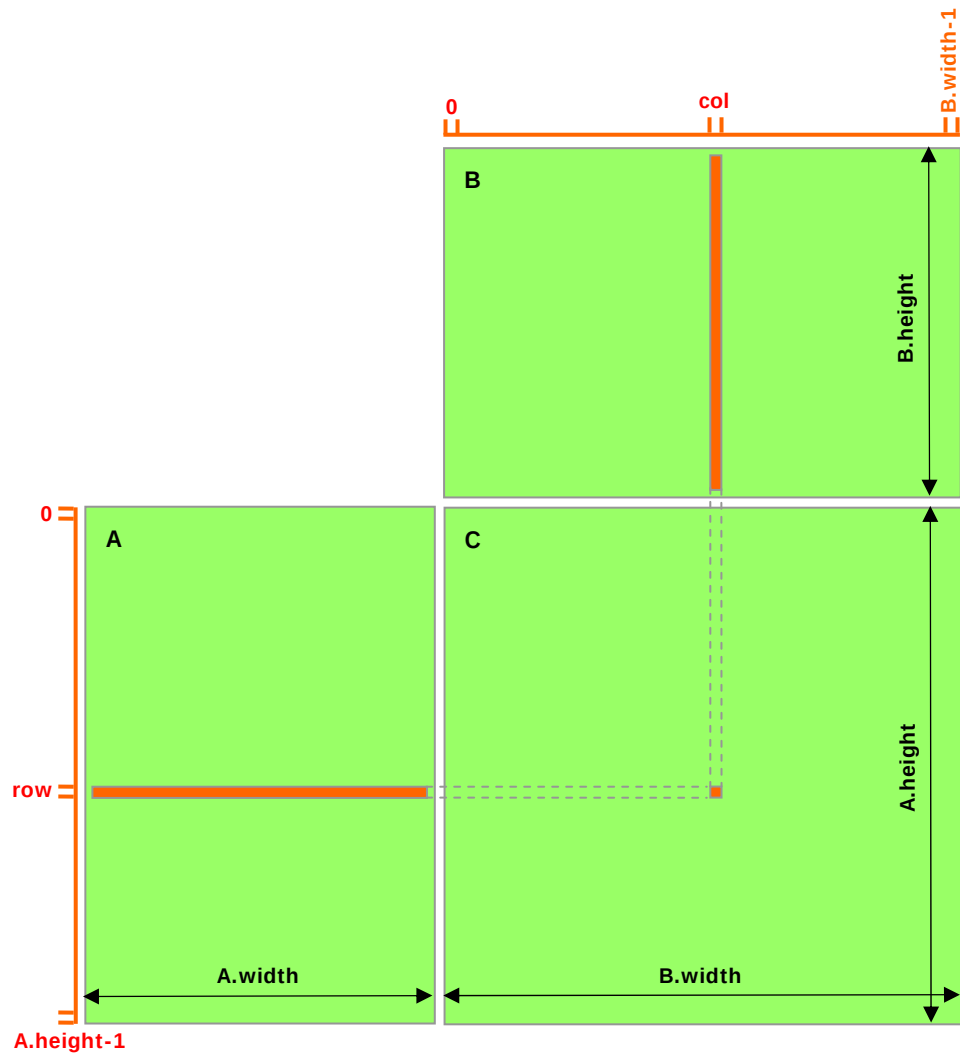


Figure 3-1. Matrix Multipliation without Shared Memory

The following code sample is an implementation of matrix multiplication that does take advantage of shared memory. In this implementation, each thread block is responsible for computing one square sub-matrix C_{sub} of C and each thread within the block is responsible for computing one element of C_{sub} . As illustrated in Figure 3-2, C_{sub} is equal to the product of two rectangular matrices: the sub-matrix of A of dimension $(A.width, block_size)$ that has the same line indices as C_{sub} , and the sub-matrix of B of dimension $(block_size, A.width)$ that has the same column indices as C_{sub} . In order to fit into the device's resources, these two rectangular matrices are divided into as many square matrices of dimension $block_size$ as necessary and C_{sub} is computed as the sum of the products of these square matrices. Each of these products is performed by first loading the two corresponding square matrices from global memory to shared memory with one thread loading one element of each matrix, and then by having each thread compute one element of the product. Each thread accumulates the result of each of these products into a register and once done writes the result to global memory.

By blocking the computation this way, we take advantage of fast shared memory and save a lot of global memory bandwidth since A is only read $(B.width / block_size)$ times from global memory and B is read $(A.height / block_size)$ times.

The *Matrix* type from the previous code sample is augmented with a *stride* field, so that sub-matrices can be efficiently represented with the same type. `__device__` functions (see Section B.1.1) are used to get and set elements and build any sub-matrix from a matrix.

```
// Matrices are stored in row-major order:
// M(row, col) = *(M.elements + row * M.stride + col)
typedef struct {
    int width;
    int height;
    int stride;
    float* elements;
} Matrix;

// Get a matrix element
__device__ float GetElement(const Matrix A, int row, int col)
{
    return A.elements[row * A.stride + col];
}

// Set a matrix element
__device__ void SetElement(Matrix A, int row, int col,
                           float value)
{
    A.elements[row * A.stride + col] = value;
}

// Get the BLOCK_SIZExBLOCK_SIZE sub-matrix Asub of A that is
// located col sub-matrices to the right and row sub-matrices down
// from the upper-left corner of A
__device__ Matrix GetSubMatrix(Matrix A, int row, int col)
{
    Matrix Asub;
    Asub.width    = BLOCK_SIZE;
    Asub.height   = BLOCK_SIZE;
    Asub.stride   = A.stride;
    Asub.elements = &A.elements[A.stride * BLOCK_SIZE * row
                                + BLOCK_SIZE * col];

    return Asub;
}

// Thread block size
#define BLOCK_SIZE 16

// Forward declaration of the matrix multiplication kernel
__global__ void MatMulKernel(const Matrix, const Matrix, Matrix);

// Matrix multiplication - Host code
// Matrix dimensions are assumed to be multiples of BLOCK_SIZE
void MatMul(const Matrix A, const Matrix B, Matrix C)
{
    // Load A and B to device memory
    Matrix d_A;
```

```

d_A.width = d_A.stride = A.width; d_A.height = A.height;
size_t size = A.width * A.height * sizeof(float);
cudaMalloc((void**)&d_A.elements, size);
cudaMemcpy(d_A.elements, A.elements, size,
           cudaMemcpyHostToDevice);
Matrix d_B;
d_B.width = d_B.stride = B.width; d_B.height = B.height;
size = B.width * B.height * sizeof(float);
cudaMalloc((void**)&d_B.elements, size);
cudaMemcpy(d_B.elements, B.elements, size,
           cudaMemcpyHostToDevice);

// Allocate C in device memory
Matrix d_C;
d_C.width = d_C.stride = C.width; d_C.height = C.height;
size = C.width * C.height * sizeof(float);
cudaMalloc((void**)&d_C.elements, size);

// Invoke kernel
dim3 dimBlock(BLOCK_SIZE, BLOCK_SIZE);
dim3 dimGrid(B.width / dimBlock.x, A.height / dimBlock.y);
MatMulKernel<<<dimGrid, dimBlock>>>(d_A, d_B, d_C);

// Read C from device memory
cudaMemcpy(C.elements, d_C.elements, size,
           cudaMemcpyDeviceToHost);

// Free device memory
cudaFree(d_A.elements);
cudaFree(d_B.elements);
cudaFree(d_C.elements);
}

// Matrix multiplication kernel called by MatMul()
__global__ void MatMulKernel(Matrix A, Matrix B, Matrix C)
{
    // Block row and column
    int blockRow = blockIdx.y;
    int blockCol = blockIdx.x;

    // Each thread block computes one sub-matrix Csub of C
    Matrix Csub = GetSubMatrix(C, blockRow, blockCol);

    // Each thread computes one element of Csub
    // by accumulating results into Cvalue
    float Cvalue = 0;

    // Thread row and column within Csub
    int row = threadIdx.y;
    int col = threadIdx.x;

    // Loop over all the sub-matrices of A and B that are
    // required to compute Csub
    // Multiply each pair of sub-matrices together
    // and accumulate the results
    for (int m = 0; m < (A.width / BLOCK_SIZE); ++m) {

```

```

// Get sub-matrix Asub of A
Matrix Asub = GetSubMatrix(A, blockRow, m);

// Get sub-matrix Bsub of B
Matrix Bsub = GetSubMatrix(B, m, blockCol);

// Shared memory used to store Asub and Bsub respectively
__shared__ float As[BLOCK_SIZE][BLOCK_SIZE];
__shared__ float Bs[BLOCK_SIZE][BLOCK_SIZE];

// Load Asub and Bsub from device memory to shared memory
// Each thread loads one element of each sub-matrix
As[row][col] = GetElement(Asub, row, col);
Bs[row][col] = GetElement(Bsub, row, col);

// Synchronize to make sure the sub-matrices are loaded
// before starting the computation
__syncthreads();

// Multiply Asub and Bsub together
for (int e = 0; e < BLOCK_SIZE; ++e)
    Cvalue += As[row][e] * Bs[e][col];

// Synchronize to make sure that the preceding
// computation is done before loading two new
// sub-matrices of A and B in the next iteration
__syncthreads();
}

// Write Csub to device memory
// Each thread writes one element
SetElement(Csub, row, col, Cvalue);
}

```

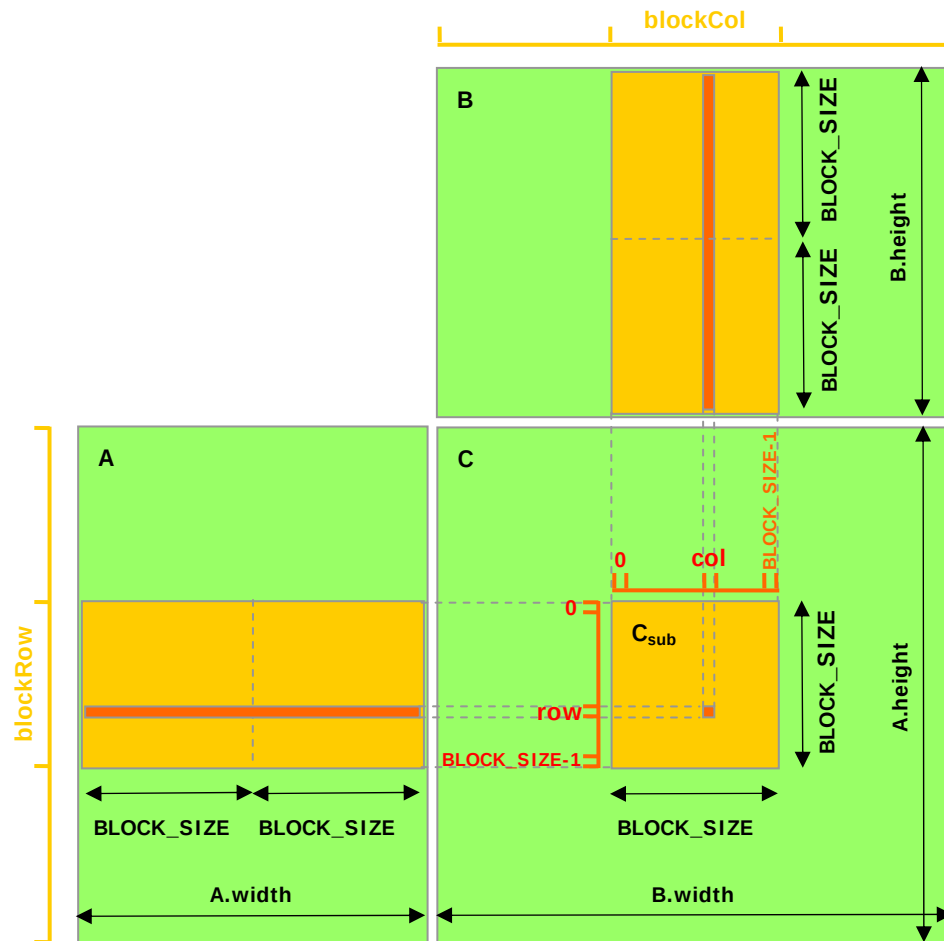


Figure 3-2. Matrix Multipliation with Shared Memory

3.2.3 Multiple Devices

A host system can have multiple devices. These devices can be enumerated, their properties can be queried, and one of them can be selected for kernel executions.

Several host threads can execute device code on the same device, but by design, a host thread can execute device code on only one device at any given time. As a consequence, multiple host threads are required to execute device code on multiple devices. Also, any CUDA resources created through the runtime in one host thread cannot be used by the runtime from another host thread.

The following code sample enumerates all devices in the system and retrieves their properties. It also determines the number of CUDA-enabled devices.

```
int deviceCount;
cudaGetDeviceCount(&deviceCount);
int device;
for (device = 0; device < deviceCount; ++device) {
    cudaDeviceProp deviceProp;
    cudaGetDeviceProperties(&deviceProp, device);
    if (dev == 0) {
```

```

    if (deviceProp.major == 9999 && deviceProp.minor == 9999)
        printf("There is no device supporting CUDA.\n");
    else if (deviceCount == 1)
        printf("There is 1 device supporting CUDA\n");
    else
        printf("There are %d devices supporting CUDA\n",
               deviceCount);
}
}

```

By default, the device associated to the host thread is implicitly selected as device 0 as soon as a non-device management runtime function is called (see Section 3.5 for exceptions). Any other device can be selected by calling **cudaSetDevice()** first. After a device has been selected, either implicitly or explicitly, any subsequent explicit call to **cudaSetDevice()** will fail up until **cudaThreadExit()** is called. **cudaThreadExit()** cleans up all runtime-related resources associated with the calling host thread. Any subsequent API call reinitializes the runtime.

3.2.4 Texture Memory

CUDA supports a subset of the texturing hardware that the GPU uses for graphics to access texture memory. Reading data from texture memory instead of global memory can have several performance benefits as described in Section 5.1.2.4.

Texture memory is read from kernels using device functions called *texture fetches*, described in Section B.8. The first parameter of a texture fetch specifies an object called a *texture reference*.

A texture reference defines which part of texture memory is fetched. As detailed in Section 3.2.4.3, it must be bound through runtime functions to some region of memory, called a *texture*, before it can be used by a kernel. Several distinct texture references might be bound to the same texture or to textures that overlap in memory.

A texture reference has several attributes. One of them is its dimensionality that specifies whether the texture is addressed as a one-dimensional array using one *texture coordinate*, a two-dimensional array using two texture coordinates, or a three-dimensional array using three texture coordinates. Elements of the array are called *texels*, short for “texture elements.”

Other attributes define the input and output data types of the texture fetch, as well as how the input coordinates are interpreted and what processing should be done.

A texture can be any region of linear memory or a CUDA array.

CUDA arrays are opaque memory layouts optimized for texture fetching. They are one-dimensional, two-dimensional, or three-dimensional and composed of elements, each of which has 1, 2 or 4 components that may be signed or unsigned 8-, 16- or 32-bit integers, 16-bit floats (currently only supported through the driver API), or 32-bit floats. CUDA arrays are only readable by kernels through texture fetching and may only be bound to texture references with the same number of packed components.

3.2.4.1 Texture Reference Declaration

Some of the attributes of a texture reference are immutable and must be known at compile time; they are specified when declaring the texture reference. A texture reference is declared at file scope as a variable of type **texture**:

```
texture<Type, Dim, ReadMode> texRef;
```

where:

- **Type** specifies the type of data that is returned when fetching the texture; **Type** is restricted to the basic integer and single-precision floating-point types and any of the 1-, 2-, and 4-component vector types defined in Section B.3.1;
- **Dim** specifies the dimensionality of the texture reference and is equal to 1, 2, or 3; **Dim** is an optional argument which defaults to 1;
- **ReadMode** is equal to **cudaReadModeNormalizedFloat** or **cudaReadModeElementType**; if it is **cudaReadModeNormalizedFloat** and **Type** is a 16-bit or 8-bit integer type, the value is actually returned as floating-point type and the full range of the integer type is mapped to [0.0, 1.0] for unsigned integer type and [-1.0, 1.0] for signed integer type; for example, an unsigned 8-bit texture element with the value 0xff reads as 1; if it is **cudaReadModeElementType**, no conversion is performed; **ReadMode** is an optional argument which defaults to **cudaReadModeElementType**.

3.2.4.2 Runtime Texture Reference Attributes

The other attributes of a texture reference are mutable and can be changed at runtime through the host runtime. They specify whether texture coordinates are normalized or not, the addressing mode, and texture filtering, as detailed below.

By default, textures are referenced using floating-point coordinates in the range [0, N) where N is the size of the texture in the dimension corresponding to the coordinate. For example, a texture that is 64×32 in size will be referenced with coordinates in the range [0, 63] and [0, 31] for the x and y dimensions, respectively. Normalized texture coordinates cause the coordinates to be specified in the range [0.0, 1.0) instead of [0, N), so the same 64×32 texture would be addressed by normalized coordinates in the range [0, 1) in both the x and y dimensions. Normalized texture coordinates are a natural fit to some applications' requirements, if it is preferable for the texture coordinates to be independent of the texture size.

The addressing mode defines what happens when texture coordinates are out of range. When using unnormalized texture coordinates, texture coordinates outside the range [0, N) are clamped: Values below 0 are set to 0 and values greater or equal to N are set to N-1. Clamping is also the default addressing mode when using normalized texture coordinates: Values below 0.0 or above 1.0 are clamped to the range [0.0, 1.0). For normalized coordinates, the “wrap” addressing mode also may be specified. Wrap addressing is usually used when the texture contains a periodic signal. It uses only the fractional part of the texture coordinate; for example, 1.25 is treated the same as 0.25 and -1.25 is treated the same as 0.75.

Linear texture filtering may be done only for textures that are configured to return floating-point data. It performs low-precision interpolation between neighboring texels. When enabled, the texels surrounding a texture fetch location are read and the return value of the texture fetch is interpolated based on where the texture coordinates fell between the texels. Simple linear interpolation is performed for one-

dimensional textures and bilinear interpolation is performed for two-dimensional textures.

Appendix E gives more details on texture fetching.

3.2.4.3 Texture Binding

As explained in the reference manual, the runtime API has a *low-level* C-style interface and a *high-level* C++-style interface. The **texture** type is defined in the high-level API as a structure publicly derived from the **textureReference** type defined in the low-level API as such:

```
struct textureReference {
    int normalized;
    enum cudaTextureFilterMode filterMode;
    enum cudaTextureAddressMode addressMode[3];
    struct cudaChannelFormatDesc channelDesc;
}
```

- **normalized** specifies whether texture coordinates are normalized or not; if it is non-zero, all elements in the texture are addressed with texture coordinates in the range **[0,1]** rather than in the range **[0,width-1]**, **[0,height-1]**, or **[0,depth-1]** where **width**, **height**, and **depth** are the texture sizes;
- **filterMode** specifies the filtering mode, that is how the value returned when fetching the texture is computed based on the input texture coordinates; **filterMode** is equal to **cudaFilterModePoint** or **cudaFilterModeLinear**; if it is **cudaFilterModePoint**, the returned value is the texel whose texture coordinates are the closest to the input texture coordinates; if it is **cudaFilterModeLinear**, the returned value is the linear interpolation of the two (for a one-dimensional texture), four (for a two-dimensional texture), or eight (for a three-dimensional texture) texels whose texture coordinates are the closest to the input texture coordinates; **cudaFilterModeLinear** is only valid for returned values of floating-point type;
- **addressMode** specifies the addressing mode, that is how out-of-range texture coordinates are handled; **addressMode** is an array of size three whose first, second, and third elements specify the addressing mode for the first, second, and third texture coordinates, respectively; the addressing mode is equal to either **cudaAddressModeClamp**, in which case out-of-range texture coordinates are clamped to the valid range, or **cudaAddressModeWrap**, in which case out-of-range texture coordinates are wrapped to the valid range; **cudaAddressModeWrap** is only supported for normalized texture coordinates;
- **channelDesc** describes the format of the value that is returned when fetching the texture; **channelDesc** is of the following type:

```
struct cudaChannelFormatDesc {
    int x, y, z, w;
    enum cudaChannelFormatKind f;
};
```

where **x**, **y**, **z**, and **w** are equal to the number of bits of each component of the returned value and **f** is:

- **cudaChannelFormatKindSigned** if these components are of signed integer type,

- **cudaChannelFormatKindUnsigned** if they are of unsigned integer type,
- **cudaChannelFormatKindFloat** if they are of floating point type.

normalized, **addressMode**, and **filterMode** may be directly modified in host code.

Before a kernel can use a texture reference to read from texture memory, the texture reference must be bound to a texture using **cudaBindTexture()** or **cudaBindTextureToArray()**. **cudaUnbindTexture()** is used to unbind a texture reference.

The following code samples bind a texture reference to linear memory pointed to by **devPtr**:

- ❑ Using the low-level API:

```
texture<float, 2, cudaReadModeElementType> texRef;
textureReference* texRefPtr;
cudaGetTextureReference(&texRefPtr, "texRef");
cudaChannelFormatDesc channelDesc =
    cudaCreateChannelDesc<float>();
cudaBindTexture2D(0, texRefPtr, devPtr, &channelDesc,
    width, height, pitch);
```

- ❑ Using the high-level API:

```
texture<float, 2, cudaReadModeElementType> texRef;
cudaChannelFormatDesc channelDesc =
    cudaCreateChannelDesc<float>();
cudaBindTexture2D(0, texRef, devPtr, &channelDesc,
    width, height, pitch);
```

The following code samples bind a texture reference to a CUDA array **cuArray**:

- ❑ Using the low-level API:

```
texture<float, 2, cudaReadModeElementType> texRef;
textureReference* texRefPtr;
cudaGetTextureReference(&texRefPtr, "texRef");
cudaChannelFormatDesc channelDesc;
cudaGetChannelDesc(&channelDesc, cuArray);
cudaBindTextureToArray(texRef, cuArray, &channelDesc);
```

- ❑ Using the high-level API:

```
texture<float, 2, cudaReadModeElementType> texRef;
cudaBindTextureToArray(texRef, cuArray);
```

The format specified when binding a texture to a texture reference must match the parameters specified when declaring the texture reference; otherwise, the results of texture fetches are undefined.

The following code sample applies some simple transformation kernel to a

```
// 2D float texture
texture<float, 2, cudaReadModeElementType> texRef;

// Simple transformation kernel
__global__ void transformKernel(float* output,
                                int width, int height, float theta)
{
    // Calculate normalized texture coordinates
    unsigned int x = blockIdx.x * blockDim.x + threadIdx.x;
```

```

    unsigned int y = blockIdx.y * blockDim.y + threadIdx.y;

    float u = x / (float)width;
    float v = y / (float)height;

    // Transform coordinates
    u -= 0.5f;
    v -= 0.5f;
    float tu = u * cosf(theta) - v * sinf(theta) + 0.5f;
    float tv = v * cosf(theta) + u * sinf(theta) + 0.5f;

    // Read from texture and write to global memory
    output[y * width + x] = tex2D(tex, tu, tv);
}

// Host code
int main()
{
    // Allocate CUDA array in device memory
    cudaChannelFormatDesc channelDesc =
        cudaCreateChannelDesc(32, 0, 0, 0,
                               cudaChannelFormatKindFloat);

    cudaArray* cuArray;
    cudaMallocArray(&cuArray, &channelDesc, width, height);

    // Copy to device memory some data located at address h_data
    // in host memory
    cudaMemcpyToArray(cuArray, 0, 0, h_data, size,
                     cudaMemcpyHostToDevice);

    // Set texture parameters
    texRef.addressMode[0] = cudaAddressModeWrap;
    texRef.addressMode[1] = cudaAddressModeWrap;
    texRef.filterMode      = cudaFilterModeLinear;
    texRef.normalized      = true;

    // Bind the array to the texture
    cudaBindTextureToArray(texRef, cuArray, channelDesc);

    // Allocate result of transformation in device memory
    float* output;
    cudaMalloc((void*)&output, width * height * sizeof(float));

    // Invoke kernel
    dim3 dimBlock(16, 16);
    dim3 dimGrid((width + dimBlock.x - 1) / dimBlock.x,
                 (height + dimBlock.y - 1) / dimBlock.y);
    transformKernel<<<dimGrid, dimBlock>>>(output, width, height,
                                           angle);

    // Free device memory
    cudaFreeArray(cuArray);
    cudaFree(output);
}

```

3.2.5 Page-Locked Host Memory

The runtime also provides functions to allocate and free *page-locked* (also known as *pinned*) host memory – as opposed to regular pageable host memory allocated by **malloc()**: **cudaHostAlloc()** and **cudaFreeHost()**.

Using page-locked host memory has several benefits:

- ❑ Bandwidth between host memory and device memory is higher if host memory is allocated as page-locked and even higher if in addition it is allocated as write-combining as described in Section 3.2.5.2;
- ❑ Copies between page-locked host memory and device memory can be performed concurrently with kernel execution for some devices as mentioned in Section 3.2.6;
- ❑ On some devices, page-locked host memory can be mapped into the device's address space, eliminating the need to copy it to or from device memory as detailed in Section 3.2.5.3.

Page-locked host memory is a scarce resource however, so allocations in page-locked memory will start failing long before allocations in pageable memory. In addition, by reducing the amount of physical memory available to the operating system for paging, allocating too much page-locked memory reduces overall system performance.

The simple zero-copy SDK sample comes with a detailed document on the page-locked memory APIs.

3.2.5.1 Portable Memory

A block of page-locked memory can be used by any host threads, but by default, the benefits of using page-locked memory described above are only available for the thread that allocates it. To make these advantages available to all threads, it needs to be allocated by passing flag **cudaHostAllocPortable** to **cudaHostAlloc()**.

3.2.5.2 Write-Combining Memory

By default page-locked host memory is allocated as cacheable. It can optionally be allocated as *write-combining* instead by passing flag **cudaHostAllocWriteCombined** to **cudaHostAlloc()**. Write-combining memory frees up L1 and L2 cache resources, making more cache available to the rest of the application. In addition, write-combining memory is not snooped during transfers across the PCI Express bus, which can improve transfer performance by up to 40%.

Reading from write-combining memory from the host is prohibitively slow, so write-combining memory should in general be used for memory that the host only writes to.

3.2.5.3 Mapped Memory

On some devices, a block of page-locked host memory can also be mapped into the device's address space by passing flag **cudaHostAllocMapped** to **cudaHostAlloc()**. Such a block has therefore two addresses: one in host memory and one in device memory. The host memory pointer is returned by **cudaHostAlloc()** and the device memory pointer can be retrieved using

cudaHostGetDevicePointer() and used to access the block from within a kernel.

Accessing host memory directly from within a kernel has several advantages:

- ❑ There is no need to allocate a block in device memory and copy data between this block and the block in host memory; data transfers are implicitly performed as needed by the kernel;
- ❑ There is no need to use streams (see Section 3.2.6.1) to overlap data transfers with kernel execution; the kernel-originated data transfers automatically overlap with kernel execution.

Since mapped page-locked memory is shared between host and device however, the application must synchronize memory accesses using streams or events (see Section 3.2.6) to avoid any potential read-after-write, write-after-read, or write-after-write hazards.

A block of page-locked host memory can be allocated as both mapped and portable (see Section 3.2.5.1), in which case each host thread that needs to map the block to its device address space must call **cudaHostGetDevicePointer()** to retrieve a device pointer, as device pointers will generally differ from one host thread to the other.

To be able to retrieve the device pointer to any mapped page-locked memory within a given host thread, page-locked memory mapping must be enabled by calling **cudaSetDeviceFlags()** with the **cudaDeviceMapHost** flag before any other CUDA calls is performed by the thread. Otherwise, **cudaHostGetDevicePointer()** will return an error.

cudaHostGetDevicePointer() also returns an error if the device does not support mapped page-locked host memory.

Applications may query whether a device supports mapped page-locked host memory or not by calling **cudaGetDeviceProperties()** and checking the **canMapHostMemory** property.

3.2.6 Asynchronous Concurrent Execution

In order to facilitate concurrent execution between host and device, some functions are asynchronous: Control is returned to the host thread before the device has completed the requested task. These are:

- ❑ Kernel launches;
- ❑ The functions that perform memory copies and are suffixed with **Async**;
- ❑ The functions that perform device ↔ device memory copies;
- ❑ The functions that set memory.

Some devices can also perform copies between page-locked host memory and device memory concurrently with kernel execution. Applications may query this capability by calling **cudaGetDeviceProperties()** and checking the **deviceOverlap** property. This capability is currently supported only for memory copies that do not involve CUDA arrays or 2D arrays allocated through **cudaMallocPitch()** (see Section 3.2.1).

3.2.6.1 Stream

Applications manage concurrency through *streams*. A stream is a sequence of commands that execute in order. Different streams, on the other hand, may execute their commands out of order with respect to one another or concurrently.

A stream is defined by creating a stream object and specifying it as the stream parameter to a sequence of kernel launches and host ↔ device memory copies. The following code sample creates two streams and allocates an array **hostPtr** of **float** in page-locked memory.

```
cudaStream_t stream[2];
for (int i = 0; i < 2; ++i)
    cudaStreamCreate(&stream[i]);
float* hostPtr;
cudaMallocHost((void**)&hostPtr, 2 * size);
```

Each of these streams is defined by the following code sample as a sequence of one memory copy from host to device, one kernel launch, and one memory copy from device to host:

```
for (int i = 0; i < 2; ++i)
    cudaMemcpyAsync(inputDevPtr + i * size, hostPtr + i * size,
                    size, cudaMemcpyHostToDevice, stream[i]);
for (int i = 0; i < 2; ++i)
    myKernel<<<100, 512, 0, stream[i]>>>
        (outputDevPtr + i * size, inputDevPtr + i * size, size);
for (int i = 0; i < 2; ++i)
    cudaMemcpyAsync(hostPtr + i * size, outputDevPtr + i * size,
                    size, cudaMemcpyDeviceToHost, stream[i]);
cudaThreadSynchronize();
```

Each stream copies its portion of input array **hostPtr** to array **inputDevPtr** in device memory, processes **inputDevPtr** on the device by calling **myKernel()**, and copies the result **outputDevPtr** back to the same portion of **hostPtr**. Processing **hostPtr** using two streams allows for the memory copies of one stream to overlap with the kernel execution of the other stream. **hostPtr** must point to page-locked host memory for any overlap to occur. **cudaThreadSynchronize()** is called in the end to make sure all streams are finished before proceeding further.

cudaStreamSynchronize() can be used to synchronize the host with a specific stream, allowing other streams to continue executing on the device. Streams are released by calling **cudaStreamDestroy()**.

```
for (int i = 0; i < 2; ++i)
    cudaStreamDestroy(&stream[i]);
```

Any kernel launch, memory set, or memory copy function without a stream parameter or with a zero stream parameter begins only after all preceding commands are done, including commands that are part of streams, and no subsequent command may begin until it is done.

cudaStreamQuery() provides applications with a way to know if all preceding commands in a stream have completed. **cudaStreamSynchronize()** provides a way to explicitly force the runtime to wait until all preceding commands in a stream have completed.

Similarly, with **cudaThreadSynchronize()** forces the runtime to wait until all preceding device tasks in all streams have completed. To avoid unnecessary slowdowns, these functions are best used for timing purposes or to isolate a launch

or memory copy that is failing. **cudaStreamDestroy()** waits for all preceding tasks in the given stream to complete before destroying the stream and returning control to the host thread.

Two commands from different streams cannot run concurrently if either a page-locked host memory allocation, a device memory allocation, a device memory set, a device $\leftrightarrow$ device memory copy, or any CUDA command to stream 0 is called in-between them by the host thread.

Programmers can globally disable asynchronous execution for all CUDA applications running on a system by setting the **CUDA_LAUNCH_BLOCKING** environment variable to 1. This feature is provided for debugging purposes only and should never be used as a way to make production software run reliably.

3.2.6.2 Event

The runtime also provides a way to closely monitor the device's progress, as well as perform accurate timing, by letting the application asynchronously record *events* at any point in the program and query when these events are actually recorded. An event is recorded when all tasks – or optionally, all commands in a given stream – preceding the event have completed. Events in stream zero are recorded after all preceding tasks/commands from all streams are completed by the device.

The following code sample creates two events:

```
cudaEvent_t start, stop;
cudaEventCreate(&start);
cudaEventCreate(&stop);
```

These events can be used to time the code sample of the previous section the following way:

```
cudaEventRecord(start, 0);
for (int i = 0; i < 2; ++i)
    cudaMemcpyAsync(inputDev + i * size, inputHost + i * size,
                    size, cudaMemcpyHostToDevice, stream[i]);
for (int i = 0; i < 2; ++i)
    myKernel<<<100, 512, 0, stream[i]>>>
        (outputDev + i * size, inputDev + i * size, size);
for (int i = 0; i < 2; ++i)
    cudaMemcpyAsync(outputHost + i * size, outputDev + i * size,
                    size, cudaMemcpyDeviceToHost, stream[i]);
cudaEventRecord(stop, 0);
cudaEventSynchronize(stop);
float elapsedTime;
cudaEventElapsedTime(&elapsedTime, start, stop);
```

They are destroyed this way:

```
cudaEventDestroy(start);
cudaEventDestroy(stop);
```

3.2.6.3 Synchronous Calls

When a synchronous function is called, control is not returned to the host thread before the device has completed the requested task. Whether the host thread will then yield, block, or spin can be specified by calling **cudaSetDeviceFlags()** with some specific flags (see reference manual for details) before any other CUDA calls is performed by the host thread.

3.2.7 OpenGL Interoperability

OpenGL buffer objects may be mapped into the address space of CUDA, either to enable CUDA to read data written by OpenGL or to enable CUDA to write data for consumption by OpenGL.

Interoperability with OpenGL requires that the CUDA device be specified by **cudaGLSetGLDevice()** before any other runtime calls.

A buffer object must be registered to CUDA before it can be mapped. This is done with **cudaGLRegisterBufferObject()**:

```
GLuint bufferObj;
cudaGLRegisterBufferObject(bufferObj);
```

Once it is registered, a buffer object can be read from or written to by kernels using the device memory address returned by **cudaGLMapBufferObject()**:

```
GLuint bufferObj;
float* devPtr;
cudaGLMapBufferObject((void*)&devPtr, bufferObj);
```

Unmapping is done with **cudaGLUnmapBufferObject()** and unregistering with **cudaGLUnregisterBufferObject()**.

The following code sample uses a kernel to dynamically modify a 2D **width x height** grid of vertices stored in a vertex buffer object:

```
GLuint positionsVB0;
int main()
{
    // Explicitly set device
    cudaGLSetGLDevice(0);

    // Initialize OpenGL and GLUT
    ...
    glutDisplayFunc(display);

    // Create buffer object and register it with CUDA
    glGenBuffers(1, positionsVB0);
    glBindBuffer(GL_ARRAY_BUFFER, &vb0);
    unsigned int size = width * height * 4 * sizeof(float);
    glBufferData(GL_ARRAY_BUFFER, size, 0, GL_DYNAMIC_DRAW);
    glBindBuffer(GL_ARRAY_BUFFER, 0);
    cudaGLRegisterBufferObject(positionsVB0);

    // Launch rendering loop
    glutMainLoop();
}

void display()
{
    // Map buffer object for writing from CUDA
    float4* positions;
    cudaGLMapBufferObject((void*)&positions, positionsVB0);

    // Execute kernel
    dim3 dimBlock(16, 16, 1);
    dim3 dimGrid(width / dimBlock.x, height / dimBlock.y, 1);
    createVertices<<<dimGrid, dimBlock>>>(positions, time,
```

```

width, height);

// Unmap buffer object
cudaGLUnmapBufferObject(positionsVBO);

// Render from buffer object
glClear(GL_COLOR_BUFFER_BIT | GL_DEPTH_BUFFER_BIT);
glBindBuffer(GL_ARRAY_BUFFER, positionsVBO);
glVertexPointer(4, GL_FLOAT, 0, 0);
glEnableClientState(GL_VERTEX_ARRAY);
glDrawArrays(GL_POINTS, 0, width * height);
glDisableClientState(GL_VERTEX_ARRAY);

// Swap buffers
glutSwapBuffers();
glutPostRedisplay();
}

void deleteVBO()
{
    cudaGLUnregisterBufferObject(positionsVBO);
    glDeleteBuffers(1, &positionsVBO);
}

__global__ void createVertices(float4* positions, float time,
                               unsigned int width, unsigned int height)
{
    unsigned int x = blockIdx.x * blockDim.x + threadIdx.x;
    unsigned int y = blockIdx.y * blockDim.y + threadIdx.y;

    // Calculate uv coordinates
    float u = x / (float)width;
    float v = y / (float)height;
    u = u * 2.0f - 1.0f;
    v = v * 2.0f - 1.0f;

    // calculate simple sine wave pattern
    float freq = 4.0f;
    float w = sinf(u * freq + time)
              * cosf(v * freq + time) * 0.5f;

    // Write positions
    positions[y * width + x] = make_float4(u, w, v, 1.0f);
}

```

On Windows and for Quadro GPUs, **cudaWGLGetDevice()** can be used to retrieve the CUDA device associated to the handle returned by **WGL_NV_gpu_affinity()**. Quadro GPUs offer higher performance OpenGL interoperability than GeForce and Tesla GPUs in a multi-GPU configuration where OpenGL rendering is performed on the Quadro GPU and CUDA computations are performed on other GPUs in the system.

3.2.8 Direct3D Interoperability

Direct3D resources may be mapped into the address space of CUDA, either to enable CUDA to read data written by Direct3D or to enable CUDA to write data for consumption by Direct3D.

Direct3D interoperability is supported for Direct3D 9.0 and Direct3D 10.0.

There are restrictions on which resources can be mapped as detailed in the reference manual for **cudaD3D9RegisterResource()** (resp. **cudaD3D10RegisterResource()**).

A CUDA context may interoperate with only one Direct3D device at a time and the CUDA context and Direct3D device must be created on the same GPU. Moreover, the Direct3D device must be created with the **D3DCREATE_HARDWARE_VERTEXPROCESSING** flag.

Interoperability with Direct3D requires that the Direct3D device be specified by **cudaD3D9SetDirect3DDevice()** (resp. **cudaD3D10SetDirect3DDevice()**) before any other runtime calls.

Direct3D resources can then be registered to CUDA using **cudaD3D9RegisterResource()** (resp. **cudaD3D10RegisterResource()**):

```
IDirect3DVertexBuffer9* buffer;
cudaD3D9RegisterResource(buffer, cudaD3D9RegisterFlagsNone);
IDirect3DSurface9* surface;
cudaD3D9RegisterResource(surface, cudaD3D9RegisterFlagsNone);
```

```
ID3D10Buffer* buffer;
cudaD3D10RegisterResource(buffer, cudaD3D10RegisterFlagsNone);
ID3D10Texture2D* tex2D;
cudaD3D10RegisterResource(tex2D, cudaD3D10RegisterFlagsNone);
```

cudaD3D9RegisterResource() (resp. **cudaD3D10RegisterResource()**) is potentially high-overhead and typically called only once per resource. Unregistering is done with **cudaD3D9UnregisterVertexBuffer()** (resp. **cudaD3D10UnregisterVertexBuffer()**).

Once a resource is registered to CUDA, it can be mapped and unmapped as many times as necessary using **cudaD3D9MapResources()** (resp. **cudaD3D10MapResources()**) and **cudaD3D9UnmapResources()** (resp. **cudaD3D10UnmapResources()**).

A mapped resource can be read from or written to by kernels using the device memory address returned by **cudaD3D9ResourceGetMappedPointer()** (resp. **cudaD3D10ResourceGetMappedPointer()**) and the size and pitch information returned by **cudaD3D9ResourceGetMappedSize()** (resp. **cudaD3D10ResourceGetMappedSize()**), **cudaD3D9ResourceGetMappedPitch()** (resp. **cudaD3D10ResourceGetMappedPitch()**), and **cudaD3D9ResourceGetMappedPitchSlice()** (resp. **cudaD3D10ResourceGetMappedPitchSlice()**).

When applicable, a CUDA array can also be obtained from a mapped resource using **cudaD3D9ResourceGetMappedArray()** (resp. **cudaD3D10ResourceGetMappedArray()**).

Accessing a resource through Direct3D while it is mapped produces undefined results.

The following code sample uses a kernel to dynamically modify a 2D **width x height** grid of vertices stored in a vertex buffer object:

```
IDirect3D9* D3D;
IDirect3DDevice9* device;
struct CUSTOMVERTEX {
    FLOAT x, y, z;
    DWORD color;
};
IDirect3DVertexBuffer9* positionsVB;

int main()
{
    // Initialize Direct3D
    D3D = Direct3DCreate9(D3D_SDK_VERSION);

    // Get a CUDA-enabled adapter
    unsigned int adapter = 0;
    for (; adapter < g_pD3D->GetAdapterCount(); adapter++) {
        D3DADAPTER_IDENTIFIER9 adapterId;
        g_pD3D->GetAdapterIdentifier(adapter, 0, &adapterId);
        int dev;
        cudaD3D9GetDevice(&dev, adapterId.DeviceName);
        if (cudaSuccess == cudaGetLastError())
            break;
    }

    // Create device
    ...
    D3D->CreateDevice(adapter, D3DDEVTYPE_HAL, hWnd,
                     D3DCREATE_HARDWARE_VERTEXPROCESSING,
                     &params, &device);

    // Register device with CUDA
    cudaD3D9SetDirect3DDevice(device);

    // Create vertex buffer and register it with CUDA
    unsigned int size = width * height * sizeof(CUSTOMVERTEX);
    device->CreateVertexBuffer(size, 0, D3DFVF_CUSTOMVERTEX,
                              D3DP00L_DEFAULT, &positionsVB, 0);
    cudaD3D9RegisterResource(positionsVB,
                              cudaD3D9RegisterFlagsNone);
    cudaD3D9ResourceSetMapFlags(positionsVB,
                                  cudaD3D9MapFlagsWriteDiscard);

    // Launch rendering loop
    while (...) {
        ...
        Render();
        ...
    }
}

void Render()
```

```

{
    // Map vertex buffer for writing from CUDA
    float4* positions;
    cudaD3D9MapResources(1, (IDirect3DResource9**)&positionsVB);
    cudaD3D9ResourceGetMappedPointer((void**)&positions,
                                     positionsVB, 0, 0);

    // Execute kernel
    dim3 dimBlock(16, 16, 1);
    dim3 dimGrid(width / dimBlock.x, height / dimBlock.y, 1);
    createVertices<<<dimGrid, dimBlock>>>(positions, time,
                                         width, height);

    // Unmap vertex buffer
    cudaD3D9UnmapResources(1, (IDirect3DResource9**)&positionsVB);

    // Draw and present
    ...
}

void releaseVB()
{
    cudaD3D9UnregisterResource(positionsVB);
    positionsVB->Release();
}

__global__ void createVertices(float4* positions, float time,
                              unsigned int width, unsigned int height)
{
    unsigned int x = blockIdx.x * blockDim.x + threadIdx.x;
    unsigned int y = blockIdx.y * blockDim.y + threadIdx.y;

    // Calculate uv coordinates
    float u = x / (float)width;
    float v = y / (float)height;
    u = u * 2.0f - 1.0f;
    v = v * 2.0f - 1.0f;

    // Calculate simple sine wave pattern
    float freq = 4.0f;
    float w = sinf(u * freq + time)
              * cosf(v * freq + time) * 0.5f;

    // Write positions
    positions[y * width + x] =
        make_float4(u, w, v, __int_as_float(0xff00ff00));
}

ID3D10Device* device;
struct CUSTOMVERTEX {
    FLOAT x, y, z;
    DWORD color;
};
ID3D10Buffer* positionsVB;

int main()

```

```

{
    // Get a CUDA-enabled adapter
    IDXGIFactory* factory;
    CreatedDXGIFactory(__uuidof(IDXGIFactory), (void**)&factory);
    IDXGIAdapter* adapter = 0;
    for (unsigned int i = 0; !adapter; ++i) {
        if (FAILED(factory->EnumAdapters(i, &adapter))
            break;
        int dev;
        cudaD3D10GetDevice(&dev, adapter);
        if (cudaSuccess == cudaGetLastError())
            break;
        adapter->Release();
    }
    factory->Release();

    // Create swap chain and device
    ...
    D3D10CreateDeviceAndSwapChain(adapter,
                                   D3D10_DRIVER_TYPE_HARDWARE, 0,
                                   D3D10_CREATE_DEVICE_DEBUG,
                                   D3D10_SDK_VERSION,
                                   &swapChainDesc &swapChain,
                                   &device);

    adapter->Release();

    // Register device with CUDA
    cudaD3D10SetDirect3DDevice(device);

    // Create vertex buffer and register it with CUDA
    unsigned int size = width * height * sizeof(CUSTOMVERTEX);
    D3D10_BUFFER_DESC bufferDesc;
    bufferDesc.Usage          = D3D10_USAGE_DEFAULT;
    bufferDesc.ByteWidth      = size;
    bufferDesc.BindFlags      = D3D10_BIND_VERTEX_BUFFER;
    bufferDesc.CPUAccessFlags = 0;
    bufferDesc.MiscFlags      = 0;
    device->CreateBuffer(&bufferDesc, 0, &positionsVB);
    cudaD3D10RegisterResource(positionsVB,
                              cudaD3D10RegisterFlagsNone);
    cudaD3D10ResourceSetMapFlags(positionsVB,
                                  cudaD3D10MapFlagsWriteDiscard);

    // Launch rendering loop
    while (...) {
        ...
        Render();
        ...
    }
}

void Render()
{
    // Map vertex buffer for writing from CUDA
    float4* positions;
    cudaD3D10MapResources(1, (ID3D10Resource**)&positionsVB);
    cudaD3D10ResourceGetMappedPointer((void**)&positions,

```

```

                                positionsVB, 0);

    // Execute kernel
    dim3 dimBlock(16, 16, 1);
    dim3 dimGrid(width / dimBlock.x, height / dimBlock.y, 1);
    createVertices<<<dimGrid, dimBlock>>>(positions, time,
                                           width, height);

    // Unmap vertex buffer
    cudaD3D10UnmapResources(1, (ID3D10Resource**)&positionsVB);

    // Draw and present
    ...
}

void releaseVB()
{
    cudaD3D10UnregisterResource(positionsVB);
    positionsVB->Release();
}

__global__ void createVertices(float4* positions, float time,
                              unsigned int width, unsigned int height)
{
    unsigned int x = blockIdx.x * blockDim.x + threadIdx.x;
    unsigned int y = blockIdx.y * blockDim.y + threadIdx.y;

    // Calculate uv coordinates
    float u = x / (float)width;
    float v = y / (float)height;
    u = u * 2.0f - 1.0f;
    v = v * 2.0f - 1.0f;

    // Calculate simple sine wave pattern
    float freq = 4.0f;
    float w = sinf(u * freq + time)
              * cosf(v * freq + time) * 0.5f;

    // Write positions
    positions[y * width + x] =
        make_float4(u, w, v, __int_as_float(0xff00ff00));
}

```

In the following code sample, each thread accesses one pixel of a 2D surface of size **(width, height)** and pixel format **float4**:

```

// host code
void* devPtr;
cudaD3D9ResourceGetMappedPointer(&devPtr, surface, 0, 0);
size_t pitch;
cudaD3D9ResourceGetMappedPitch(&pitch, 0, surface, 0, 0);
dim3 Db = dim3(16, 16);
dim3 Dg = dim3((width+Db.x-1)/Db.x, (height+Db.y-1)/Db.y);
myKernel<<<Dg, Db>>>((unsigned char*)devPtr,
                    width, height, pitch);

// device code
__global__ void myKernel(unsigned char* surface,

```

```

        int width, int height, size_t pitch)
{
    int x = blockIdx.x * blockDim.x + threadIdx.x;
    int y = blockIdx.y * blockDim.y + threadIdx.y;
    if (x >= width || y >= height) return;
    float* pixel = (float*)(surface + y * pitch) + 4 * x;
}

// host code
void* devPtr;
cudaD3D10ResourceGetMappedPointer(&devPtr, surface, 0);
size_t pitch;
cudaD3D10ResourceGetMappedPitch(&pitch, 0, surface, 0);
dim3 Db = dim3(16, 16);
dim3 Dg = dim3((width+Db.x-1)/Db.x, (height+Db.y-1)/Db.y);
myKernel<<<Dg, Db>>>((unsigned char*)devPtr,
                    width, height, pitch);

// device code
__global__ void myKernel(unsigned char* surface,
                        int width, int height, size_t pitch)
{
    int x = blockIdx.x * blockDim.x + threadIdx.x;
    int y = blockIdx.y * blockDim.y + threadIdx.y;
    if (x >= width || y >= height) return;
    float* pixel = (float*)(surface + y * pitch) + 4 * x;
}

```

3.2.9 Error Handling

All runtime functions return an error code, but for an asynchronous function (see Section 3.2.6), this error code cannot possibly report any of the asynchronous errors that could occur on the device since the function returns before the device has completed the task; the error code only reports errors that occur on the host prior to executing the task, typically related to parameter validation; if an asynchronous error occurs, it will be reported by some subsequent unrelated runtime function call.

The only way to check for asynchronous errors just after some asynchronous function call is therefore to synchronize just after the call by calling **cudaThreadSynchronize()** (or by using any other synchronization mechanisms described in Section 3.2.6) and checking the error code returned by **cudaThreadSynchronize()**.

The runtime maintains an error variable for each host thread that is initialized to **cudaSuccess** and is overwritten by the error code every time an error occurs (be it a parameter validation error or an asynchronous error). **cudaGetLastError()** returns this variable and resets it to **cudaSuccess**.

Kernel launches do not return any error code, so **cudaGetLastError()** must be called just after the kernel launch to retrieve any pre-launch errors. To ensure that any error returned by **cudaGetLastError()** does not originate from calls prior to the kernel launch, one has to make sure that the runtime error variable is set to **cudaSuccess** just before the kernel launch, for example, by calling **cudaGetLastError()** just before the kernel launch. Kernel launches are

asynchronous, so to check for asynchronous errors, the application must synchronize in-between the kernel launch and the call to `cudaGetLastError()`.

3.2.10 Debugging using the Device Emulation Mode

CUDA-GDB can be used to debug devices of compute capability greater than 1.0. (see the CUDA-GDB user manual for supported platforms). The compiler and runtime also supports an emulation mode for the purpose of debugging that can be used even in the absence of any CUDA-enabled device. When compiling an application in this mode (using the `-deviceemu` option), the device code is compiled for and runs on the host, allowing the programmer to use the host's native debugging support to debug the application as if it were a host application. The preprocessor macro `__DEVICE_EMULATION__` is defined in this mode. All code for an application, including any libraries used, must be compiled consistently either for device emulation or for device execution. Linking code compiled for device emulation with code compiled for device execution causes the following runtime error to be returned upon initialization: `cudaErrorMixedDeviceExecution`.

When running an application in device emulation mode, the programming model is emulated by the runtime. For each thread in a thread block, the runtime creates a thread on the host. The programmer needs to make sure that:

- ❑ The host is able to run up to the maximum number of threads per block, plus one for the master thread.
- ❑ Enough memory is available to run all threads, knowing that each thread gets 256 KB of stack.

Many features provided through the device emulation mode make it a very effective debugging tool:

- ❑ By using the host's native debugging support programmers can use all features that the debugger supports, like setting breakpoints and inspecting data.
- ❑ Since device code is compiled to run on the host, the code can be augmented with code that cannot run on the device, like input and output operations to files or to the screen (`printf()`, etc.).
- ❑ Since all data resides on the host, any device- or host-specific data can be read from either device or host code; similarly, any device or host function can be called from either device or host code.
- ❑ In case of incorrect usage of the synchronization intrinsic function, the runtime detects dead lock situations.

Programmers must keep in mind that device emulation mode is emulating the device, not simulating it. Therefore, device emulation mode is very useful in finding algorithmic errors, but certain errors are hard to find:

- ❑ Race conditions are less likely to manifest themselves in device-emulation mode, since the number of threads executing simultaneously is much smaller than on an actual device.
- ❑ When dereferencing a pointer to global memory on the host or a pointer to host memory on the device, device execution almost certainly fails in some undefined way, whereas device emulation can produce correct results.
- ❑ Most of the time the same floating-point computation will not produce exactly the same result when performed on the device as when performed on the host in

device emulation mode. This is expected since in general, all you need to get different results for the same floating-point computation are slightly different compiler options, let alone different compilers, different instruction sets, or different architectures.

In particular, some host platforms store intermediate results of single-precision floating-point calculations in extended precision registers, potentially resulting in significant differences in accuracy when running in device emulation mode. When this occurs, programmers can try any of the following methods, none of which is guaranteed to work:

- Declare some floating-point variables as volatile to force single-precision storage;
- Use the **-ffloat-store** compiler option of **gcc**,
- Use the **/Op** or **/fp** compiler options of the Visual C++ compiler,
- Use **_FPU_GETCW()** and **_FPU_SETCW()** on Linux or **_controlfp()** on Windows to force single-precision floating-point computation for a portion of the code by surrounding it with

```
unsigned int originalCW;
_FPU_GETCW(originalCW);
unsigned int cw = (originalCW & ~0x300) | 0x000;
_FPU_SETCW(cw);
```

or

```
unsigned int originalCW = _controlfp(0, 0);
_controlfp(_PC_24, _MCW_PC);
```

at the beginning, to store the current value of the control word and change it to force the mantissa to be stored in 24 bits using, and with

```
_FPU_SETCW(originalCW);
```

or

```
_controlfp(originalCW, 0xffffffff);
```

at the end, to restore the original control word.

Also, for single-precision floating-point numbers, unlike compute devices (see Appendix A), host platforms usually support denormalized numbers. This can lead to dramatically different results between device emulation and device execution modes since some computation might produce a finite result in one case and an infinite result in the other.

- ❑ The warp size is equal to 1 in device emulation mode (see Section 4.1 for the definition of a warp). Therefore, the warp vote functions (described in Section B.11) produce different results than in device execution mode.

3.3 Driver API

The driver API is a handle-based, imperative API: Most objects are referenced by opaque handles that may be specified to functions to manipulate the objects.

The objects available in the driver API are summarized in Table 3-1.

Table 3-1. Objects Available in the CUDA Driver API

Object	Handle	Description
Device	CUdevice	CUDA-enabled device
Context	CUcontext	Roughly equivalent to a CPU process
Module	CUmodule	Roughly equivalent to a dynamic library
Function	CUfunction	Kernel
Heap memory	CUdeviceptr	Pointer to device memory
CUDA array	CUarray	Opaque container for one-dimensional or two-dimensional data on the device, readable via texture references
Texture reference	CUTexref	Object that describes how to interpret texture memory data

The driver API is implemented in the **nvcuda** dynamic library and all its entry points are prefixed with **cu**.

The driver API must be initialized with **cuInit()** before any function from the driver API is called. A CUDA context must then be created that is attached to a specific device and made current to the calling host thread as detailed in Section 3.3.1.

Within a CUDA context, kernels are explicitly loaded as *PTX* or binary objects by the host code as described in Section 3.3.2. Kernels written in C must therefore be compiled separately into *PTX* or binary objects. Kernels are launched using API entry points as described in Section 3.3.3.

Any application that wants to run on future device architectures must load *PTX*, not binary code. This is because binary code is architecture-specific and therefore incompatible with future architectures, whereas *PTX* code is compiled to binary code at load time by the driver.

Here is the host code of the sample from Section 2.1 written using the driver API:

```
int main()
{
    // Initialize
    if (cuInit(0) != CUDA_SUCCESS)
        exit (0);

    // Get number of devices supporting CUDA
    int deviceCount = 0;
    cuDeviceGetCount(&deviceCount);
    if (deviceCount == 0) {
        printf("There is no device supporting CUDA.\n");
        exit (0);
    }

    // Get handle for device 0
    CUdevice cuDevice = 0;
    cuDeviceGet(&cuDevice, 0);

    // Create context
    CUcontext cuContext;
    cuCtxCreate(&cuContext, 0, cuDevice);
}
```

```

// Create module from binary file
CModule cuModule;
cuModuleLoad(&cuModule, "VecAdd.ptx");

// Get function handle from module
CUfunction vecAdd;
cuModuleGetFunction(&vecAdd, cuModule, "VecAdd");

// Invoke kernel
#define ALIGN_UP(offset, alignment) \
    (offset) = ((offset) + (alignment) - 1) & ~((alignment) - 1)
int offset = 0;
void* ptr;
ptr = (void*)(size_t)A;
ALIGN_UP(offset, __alignof(ptr));
cuParamSetv(vecAdd, offset, &ptr, sizeof(ptr));
offset += sizeof(ptr);
ptr = (void*)(size_t)B;
ALIGN_UP(offset, __alignof(ptr));
cuParamSetv(vecAdd, offset, &ptr, sizeof(ptr));
offset += sizeof(ptr);
ptr = (void*)(size_t)C;
ALIGN_UP(offset, __alignof(ptr));
cuParamSetv(vecAdd, offset, &ptr, sizeof(ptr));
offset += sizeof(ptr);
cuParamSetSize(vecAdd, offset);
int threadsPerBlock = 256;
int blocksPerGrid =
    (N + threadsPerBlock - 1) / threadsPerBlock;
cuFuncSetBlockShape(vecAdd, threadsPerBlock, 1, 1);
cuLaunchGrid(vecAdd, blocksPerGrid, 1);
}

```

3.3.1 Context

A CUDA context is analogous to a CPU process. All resources and actions performed within the driver API are encapsulated inside a CUDA context, and the system automatically cleans up these resources when the context is destroyed. Besides objects such as modules and texture references, each context has its own distinct 32-bit address space. As a result, **CUdeviceptr** values from different contexts reference different memory locations.

A host thread may have only one device context current at a time. When a context is created with **cuCtxCreate()**, it is made current to the calling host thread. CUDA functions that operate in a context (most functions that do not involve device enumeration or context management) will return **CUDA_ERROR_INVALID_CONTEXT** if a valid context is not current to the thread.

Each host thread has a stack of current contexts. **cuCtxCreate()** pushes the new context onto the top of the stack. **cuCtxPopCurrent()** may be called to detach the context from the host thread. The context is then "floating" and may be pushed as the current context for any host thread. **cuCtxPopCurrent()** also restores the previous current context, if any.

A usage count is also maintained for each context. **cuCtxCreate()** creates a context with a usage count of 1. **cuCtxAttach()** increments the usage count and **cuCtxDetach()** decrements it. A context is destroyed when the usage count goes to 0 when calling **cuCtxDetach()** or **cuCtxDestroy()**.

Usage count facilitates interoperability between third party authored code operating in the same context. For example, if three libraries are loaded to use the same context, each library would call **cuCtxAttach()** to increment the usage count and **cuCtxDetach()** to decrement the usage count when the library is done using the context. For most libraries, it is expected that the application will have created a context before loading or initializing the library; that way, the application can create the context using its own heuristics, and the library simply operates on the context handed to it. Libraries that wish to create their own contexts – unbeknownst to their API clients who may or may not have created contexts of their own – would use **cuCtxPushCurrent()** and **cuCtxPopCurrent()** as illustrated in Figure 3-3.

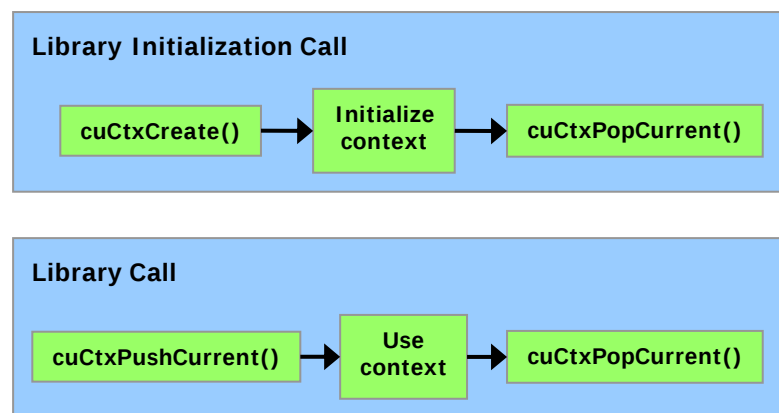


Figure 3-3. Library Context Management

3.3.2 Module

Modules are dynamically loadable packages of device code and data, akin to DLLs in Windows, that are output by **nvcc** (see Section 3.1). The names for all symbols, including functions, global variables, and texture references, are maintained at module scope so that modules written by independent third parties may interoperate in the same CUDA context.

This code sample loads a module and retrieves a handle to some kernel:

```
CUmodule cuModule;
cuModuleLoad(&cuModule, "myModule.ptx");
CUfunction myKernel;
cuModuleGetFunction(&myKernel, cuModule, "myKernel");
```

This code sample compiles and loads a new module from *PTX* code and parses compilation errors:

```
#define ERROR_BUFFER_SIZE 100
CUmodule cuModule;
CUptxas_option options[3];
void* values[3];
```

```

char* PTXCode = "some PTX code";
options[0] = CU_ASM_ERROR_LOG_BUFFER;
values[0] = (void*)malloc(ERROR_BUFFER_SIZE);
options[1] = CU_ASM_ERROR_LOG_BUFFER_SIZE_BYTES;
values[1] = (void*)ERROR_BUFFER_SIZE;
options[2] = CU_ASM_TARGET_FROM_CUCONTEXT;
values[2] = 0;
cuModuleLoadDataEx(&cuModule, PTXCode, 3, options, values);
for (int i = 0; i < values[1]; ++i) {
    // Parse error string here
}

```

3.3.3 Kernel Execution

cuFuncSetBlockShape() sets the number of threads per block for a given function, and how their threadIDs are assigned.

cuFuncSetSharedSize() sets the size of shared memory for the function.

The **cuParam*()** family of functions is used to specify the parameters that will be provided to the kernel the next time **cuLaunchGrid()** or **cuLaunch()** is invoked to launch the kernel.

The second argument of each of the **cuParam*()** functions specifies the offset of the parameter in the parameter stack. This offset must match the alignment requirement for the parameter type in device code. Alignment requirements in device code for the built-in vector types are listed in Table B-1. For all other basic types, the alignment requirement in device code matches the alignment requirement in host code and can therefore be obtained using **__alignof()**. The only exception is when the host compiler aligns **double** and **long long** (and **long** on a 64-bit system) on a one-word boundary instead of a two-word boundary (for example, using **gcc**'s compilation flag **-mno-align-double**) since in device code these types are always aligned on a two-word boundary. Also, **CUdeviceptr** is an integer, but represents a pointer, so its alignment requirement is **__alignof(void*)**. The following code sample uses a macro to adjust the offset of each parameter to meet its alignment requirement.

```

#define ALIGN_UP(offset, alignment) \
    (offset) = ((offset) + (alignment) - 1) & ~((alignment) - 1)
int offset = 0;

int i;
ALIGN_UP(offset, __alignof(i));
cuParamSeti(cuFunction, offset, i);
offset += sizeof(i);

float4 f4;
ALIGN_UP(offset, 16); // float4's alignment is 16
cuParamSetv(cuFunction, offset, &f4, sizeof(f4));
offset += sizeof(f4);

char c;
ALIGN_UP(offset, __alignof(c));
cuParamSeti(cuFunction, offset, c);
offset += sizeof(c);

```

```

float f;
ALIGN_UP(offset, __alignof(f));
cuParamSeti(cuFunction, offset, f);
offset += sizeof(f);

CUDeviceptr dptr;
// void* should be used to determine CUDeviceptr's alignment
void* ptr = (void*)(size_t)dptr;
ALIGN_UP(offset, __alignof(ptr));
cuParamSetv(cuFunction, offset, &ptr, sizeof(ptr));
offset += sizeof(ptr);

float2 f2;
ALIGN_UP(offset, 8); // float2's alignment is 8
cuParamSetv(cuFunction, offset, &f2, sizeof(f2));
offset += sizeof(f2);

cuParamSetSize(cuFunction, offset);

cuFuncSetBlockShape(cuFunction, blockDim, blockDim, 1);
cuLaunchGrid(cuFunction, gridDim, gridDim);

```

The alignment requirement of a structure is equal to the maximum of the alignment requirements of its fields. The alignment requirement of a structure that contains built-in vector types, **CUDeviceptr**, or non-aligned **double** and **long long**, might therefore differ between device code and host code. Such a structure might also be padded differently. The following structure, for example, is not padded at all in host code, but it is padded in device code with 12 bytes after field **f** since the alignment requirement for field **f4** is 16.

```

typedef struct {
    float f;
    float4 f4;
} myStruct;

```

Any parameter of type **myStruct** must therefore be passed using separate calls to **cuParam*()**, such as:

```

myStruct s;
int offset = 0;

cuParamSetv(cuFunction, offset, &s.f, sizeof(s.f));
offset += sizeof(s.f);

ALIGN_UP(offset, 16); // float4's alignment is 16
cuParamSetv(cuFunction, offset, &s.f4, sizeof(s.f4));
offset += sizeof(s.f4);

```

3.3.4 Device Memory

Linear memory is allocated using **cuMemAlloc()** or **cuMemAllocPitch()** and freed using **cuMemFree()**.

Here is the host code of the sample from Section 3.2.1 written using the driver API:

```

// Host code
int main()
{

```

```

// Initialize
if (cuInit(0) != CUDA_SUCCESS)
    exit (0);

// Get number of devices supporting CUDA
int deviceCount = 0;
cuDeviceGetCount(&deviceCount);
if (deviceCount == 0) {
    printf("There is no device supporting CUDA.\n");
    exit (0);
}

// Get handle for device 0
CUdevice cuDevice = 0;
cuDeviceGet(&cuDevice, 0);

// Create context
CUcontext cuContext;
cuCtxCreate(&cuContext, 0, cuDevice);

// Create module from binary file
CUmodule cuModule;
cuModuleLoad(&cuModule, "VecAdd.ptx");

// Get function handle from module
CUfunction vecAdd;
cuModuleGetFunction(&vecAdd, cuModule, "VecAdd");

// Allocate vectors in device memory
size_t size = N * sizeof(float);
CUdeviceptr d_A;
cuMemAlloc(&d_A, size);
CUdeviceptr d_B;
cuMemAlloc(&d_B, size);
CUdeviceptr d_C;
cuMemAlloc(&d_C, size);

// Copy vectors from host memory to device memory
// h_A and h_B are input vectors stored in host memory
cuMemcpyHtoD(d_A, h_A, size);
cuMemcpyHtoD(d_B, h_B, size);

// Invoke kernel
#define ALIGN_UP(offset, alignment) \
    (offset) = ((offset) + (alignment) - 1) & ~((alignment) - 1)
int offset = 0;
void* ptr;
ptr = (void*)(size_t)d_A;
ALIGN_UP(offset, __alignof(ptr));
cuParamSetv(vecAdd, offset, &ptr, sizeof(ptr));
offset += sizeof(ptr);
ptr = (void*)(size_t)d_B;
ALIGN_UP(offset, __alignof(ptr));
cuParamSetv(vecAdd, offset, &ptr, sizeof(ptr));
offset += sizeof(ptr);
ptr = (void*)(size_t)d_C;
ALIGN_UP(offset, __alignof(ptr));

```

```

cuParamSetv(vecAdd, offset, &ptr, sizeof(ptr));
offset += sizeof(ptr);
cuParamSetSize(VecAdd, offset);
int threadsPerBlock = 256;
int blocksPerGrid =
    (N + threadsPerBlock - 1) / threadsPerBlock;
cuFuncSetBlockShape(vecAdd, threadsPerBlock, 1, 1);
cuLaunchGrid(VecAdd, blocksPerGrid, 1);

// Copy result from device memory to host memory
// h_C contains the result in host memory
cuMemcpyDtoH(h_C, d_C, size);

// Free device memory
cuMemFree(d_A);
cuMemFree(d_B);
cuMemFree(d_C);
}

```

Linear memory can also be allocated through **cuMemAllocPitch()**. This function is recommended for allocations of 2D arrays as it makes sure that the allocation is appropriately padded to meet the alignment requirements described in Section 5.1.2.1, therefore ensuring best performance when accessing the row addresses or performing copies between 2D arrays and other regions of device memory (using the **cuMemcpy2D()**). The returned pitch (or stride) must be used to access array elements. The following code sample allocates a **width×height** 2D array of floating-point values and shows how to loop over the array elements in device code:

```

// Host code (assuming cuModule has been loaded)
CUdeviceptr devPtr;
int pitch;
cuMemAllocPitch(&devPtr, &pitch,
    width * sizeof(float), height, 4);
CUfunction myKernel;
cuModuleGetFunction(&myKernel, cuModule, "myKernel");
void* ptr = (void*)(size_t)devPtr;
cuParamSetv(myKernel, 0, &ptr, sizeof(ptr));
cuParamSetSize(myKernel, sizeof(ptr));
cuFuncSetBlockShape(myKernel, 512, 1, 1);
cuLaunchGrid(myKernel, 100, 1);

// Device code
__global__ void myKernel(float* devPtr)
{
    for (int r = 0; r < height; ++r) {
        float* row = (float*)((char*)devPtr + r * pitch);
        for (int c = 0; c < width; ++c) {
            float element = row[c];
        }
    }
}

```

The following code sample allocates a **width×height** CUDA array of one 32-bit floating-point component:

```

CUDA_ARRAY_DESCRIPTOR desc;
desc.Format = CU_AD_FORMAT_FLOAT;

```

```
desc.NumChannels = 1;
desc.Width = width;
desc.Height = height;
CUarray cuArray;
cuArrayCreate(&cuArray, &desc);
```

The reference manual lists all the various functions used to copy memory between linear memory allocated with **cuMemAlloc()**, linear memory allocated with **cuMemAllocPitch()**, and CUDA arrays.

The following code sample copies the 2D array to the CUDA array allocated in the previous code samples:

```
CUDA_MEMCPY2D copyParam;
memset(&copyParam, 0, sizeof(copyParam));
copyParam.dstMemoryType = CU_MEMORYTYPE_ARRAY;
copyParam.dstArray = cuArray;
copyParam.srcMemoryType = CU_MEMORYTYPE_DEVICE;
copyParam.srcDevice = devPtr;
copyParam.srcPitch = pitch;
copyParam.WidthInBytes = width * sizeof(float);
copyParam.Height = height;
cuMemcpy2D(&copyParam);
```

The following code sample copies some host memory array to constant memory:

```
__constant__ float constData[256];
float data[256];
CUdeviceptr devPtr;
unsigned int bytes;
cuModuleGetGlobal(&devPtr, &bytes, cuModule, "constData");
cuMemcpyHtoD(devPtr, data, bytes);
```

3.3.5 Shared Memory

The following code sample is the driver version of the host code of the sample from Section 3.2.2. Note how the **Matrix** type must be declared differently in host code since device pointers are represented as a handle of type **CUdeviceptr**.

In this sample, shared memory is statically allocated within the kernel as opposed to allocated at runtime through **cuFuncSetSharedSize()**.

```
// Matrices are stored in row-major order:
// M(row, col) = *(M.elements + row * M.stride + col)
typedef struct {
    int width;
    int height;
    int stride;
#ifdef __CUDACC__
    float* elements;
#else
    CUdeviceptr elements;
#endif
} Matrix;

// Matrix multiplication - Host code
// Matrix dimensions are assumed to be multiples of BLOCK_SIZE
void MatMul(const Matrix A, const Matrix B, Matrix C)
{
```

```

// Load A and B to device memory
Matrix d_A;
d_A.width = d_A.stride = A.width; d_A.height = A.height;
size_t size = A.width * A.height * sizeof(float);
cuMemAlloc(&d_A.elements, size);
cuMemcpyHtoD(d_A.elements, A.elements, size);
Matrix d_B;
d_B.width = d_B.stride = B.width; d_B.height = B.height;
size = B.width * B.height * sizeof(float);
cuMemAlloc(&d_B.elements, size);
cuMemcpyHtoD(d_B.elements, B.elements, size);

// Allocate C in device memory
Matrix d_C;
d_C.width = d_C.stride = C.width; d_C.height = C.height;
size = C.width * C.height * sizeof(float);
cuMemAlloc(&d_C.elements, size);

// Invoke kernel (assuming cuModule has been loaded)
CUfunction matMulKernel;
cuModuleGetFunction(&matMulKernel, cuModule, "MatMulKernel");
#define ALIGN_UP(offset, alignment) \
    (offset) = ((offset) + (alignment) - 1) & ~((alignment) - 1)
int offset = 0;
void* ptr;
ptr = (void*)(size_t)d_A;
ALIGN_UP(offset, __alignof(ptr));
cuParamSetv(matMulKernel, offset, &ptr, sizeof(ptr));
offset += sizeof(ptr);
ptr = (void*)(size_t)d_B;
ALIGN_UP(offset, __alignof(ptr));
cuParamSetv(matMulKernel, offset, &ptr, sizeof(ptr));
offset += sizeof(ptr);
ptr = (void*)(size_t)d_C;
ALIGN_UP(offset, __alignof(ptr));
cuParamSetv(matMulKernel, offset, &ptr, sizeof(ptr));
offset += sizeof(ptr);
cuParamSetSize(matMulKernel, offset);
cuFuncSetBlockShape(matMulKernel, BLOCK_SIZE, BLOCK_SIZE, 1);
cuLaunchGrid(matMulKernel,
              B.width / dimBlock.x, A.height / dimBlock.y);

// Read C from device memory
cuMemcpyDtoH(C.elements, d_C.elements, size);

// Free device memory
cuMemFree(d_A.elements);
cuMemFree(d_B.elements);
cuMemFree(d_C.elements);
}

```

3.3.6 Multiple Devices

cuDeviceGetCount() and **cuDeviceGet()** provide a way to enumerate the devices present in the system and other functions (described in the reference manual) to retrieve their properties:

```

int deviceCount;
cuDeviceGetCount(&deviceCount);
int device;
for (int device = 0; device < deviceCount; ++device) {
    CUdevice cuDevice;
    cuDeviceGet(&cuDevice, device);
    int major, minor;
    cuDeviceComputeCapability(&major, &minor, cuDevice);
}

```

3.3.7 Texture Memory

Texture binding is done using **cuTexRefSetAddress()** for linear memory and **cuTexRefSetArray()** for CUDA arrays.

If a module **cuModule** contains some texture reference **texRef** defined as

```
texture<float, 2, cudaReadModeElementType> texRef;
```

the following code sample retrieves **texRef**'s handle:

```

CUtexref cuTexRef;
cuModuleGetTexRef(&cuTexRef, cuModule, "texRef");

```

The following code sample binds **texRef** to some linear memory pointed to by **devPtr**:

```

CUDA_ARRAY_DESCRIPTOR desc;
cuTexRefSetAddress2D(cuTexRef, &desc, devPtr, pitch);

```

The following code samples bind **texRef** to a CUDA array **cuArray**:

```
cuTexRefSetArray(cuTexRef, cuArray, CU_TRSA_OVERRIDE_FORMAT);
```

The reference manual lists various functions used to set address mode, filter mode, format, and other flags for some texture reference. The format specified when binding a texture to a texture reference must match the parameters specified when declaring the texture reference; otherwise, the results of texture fetches are undefined.

The following code sample is the driver version of the host code of the sample from Section 3.2.4.3.

```

// Host code
int main()
{
    // Allocate CUDA array in device memory
    CUarray cuArray;
    CUDA_ARRAY_DESCRIPTOR desc;
    desc.Format      = CU_AD_FORMAT_FLOAT;
    desc.NumChannels = 1;
    desc.Width       = width;
    desc.Height       = height;
    cuArrayCreate(&cuArray, &desc);

    // Copy to device memory some data located at address h_data
    // in host memory
    CUDA_MEMCPY2D copyParam;
    memset(&copyParam, 0, sizeof(copyParam));
    copyParam.dstMemoryType = CU_MEMORYTYPE_ARRAY;
    copyParam.dstArray      = cuArray;
}

```

```

copyParam.srcMemoryType = CU_MEMORYTYPE_HOST;
copyParam.srcHost        = h_data;
copyParam.srcPitch        = width * sizeof(float);
copyParam.WidthInBytes    = copyParam.srcPitch;
copyParam.Height          = height;
cuMemcpy2D(&copyParam);

// Set texture parameters
CUtexref texRef;
cuModuleGetTexRef(&texRef, cuModule, "texRef");
cuTexRefSetAddressMode(texRef, 0, CU_TR_ADDRESS_MODE_WRAP);
cuTexRefSetAddressMode(texRef, 1, CU_TR_ADDRESS_MODE_WRAP);
cuTexRefSetFilterMode(texRef, CU_TR_FILTER_MODE_LINEAR);
cuTexRefSetFlags(texRef, CU_TRSF_NORMALIZED_COORDINATES);
cuTexRefSetFormat(texRef, CU_AD_FORMAT_FLOAT, 1);

// Bind the array to the texture
cuTexRefSetArray(texRef, cuArray, CU_TRSA_OVERRIDE_FORMAT);

// Allocate result of transformation in device memory
CUdeviceptr output;
cuMemAlloc((void**)&output, width * height * sizeof(float));

// Invoke kernel (assuming cuModule has been loaded)
CUfunction transformKernel;
cuModuleGetFunction(&transformKernel,
                   cuModule, "transformKernel");
#define ALIGN_UP(offset, alignment) \
    (offset) = ((offset) + (alignment) - 1) & ~((alignment) - 1)
int offset = 0;
void* ptr = (void*)(size_t)output;
ALIGN_UP(offset, __alignof(ptr));
cuParamSetv(transformKernel, offset, &ptr, sizeof(ptr));
offset += sizeof(ptr);
ALIGN_UP(offset, __alignof(width));
cuParamSeti(transformKernel, offset, width);
offset += sizeof(width);
ALIGN_UP(offset, __alignof(height));
cuParamSeti(transformKernel, offset, height);
offset += sizeof(height);
ALIGN_UP(offset, __alignof(angle));
cuParamSetf(transformKernel, offset, angle);
offset += sizeof(angle);
cuParamSetSize(transformKernel, offset);
cuParamSetTexRef(transformKernel,
                  CU_PARAM_TR_DEFAULT, texRef);
cuFuncSetBlockShape(transformKernel, 16, 16, 1);
cuLaunchGrid(transformKernel,
              (width + dimBlock.x - 1) / dimBlock.x,
              (height + dimBlock.y - 1) / dimBlock.y);

// Free device memory
cuArrayDestroy(cuArray);
cuMemFree(output);
}

```

3.3.8 Page-Locked Host Memory

Page-locked host memory can be allocated using `cuMemHostAlloc()` with optional mutually non-exclusive flags:

- ❑ **CU_MEMHOSTALLOC_PORTABLE** to allocate memory that is portable across CUDA contexts (see Section 3.2.5.1) 3.2.5.2;
- ❑ **CU_MEMHOSTALLOC_WRITECOMBINED** to allocate memory as write-combining (see Section 3.2.5.2);
- ❑ **CU_MEMHOSTALLOC_DEVICEMAP** to allocate mapped page-locked memory (see Section 3.2.5.3).

Page-locked host memory is freed using `cuMemFreeHost()`.

Page-locked memory mapping is enabled for a CUDA context by creating the context with the **CU_CTX_MAP_HOST** flag and device pointers to mapped page-locked memory are retrieved using `cuMemHostGetDevicePointer()`.

Applications may query whether a device supports mapped page-locked host memory or not by checking the

CU_DEVICE_ATTRIBUTE_CAN_MAP_HOST_MEMORY attribute using `cuDeviceGetAttribute()`.

3.3.9 Asynchronous Concurrent Execution

Applications may query if a device can perform copies between page-locked host memory and device memory concurrently with kernel execution by checking the **CU_DEVICE_ATTRIBUTE_GPU_OVERLAP** attribute using `cuDeviceGetAttribute()`.

3.3.9.1 Stream

The driver API provides functions similar to the runtime API to manage streams. The following code sample is the driver version of the code sample from Section 3.2.6.1.

```
CUstream stream[2];
for (int i = 0; i < 2; ++i)
    cuStreamCreate(&stream[i], 0);
float* hostPtr;
cuMemAllocHost((void*)&hostPtr, 2 * size);

for (int i = 0; i < 2; ++i)
    cuMemcpyHtoDAsync(inputDevPtr + i * size, hostPtr + i * size,
                      size, stream[i]);
for (int i = 0; i < 2; ++i) {
    #define ALIGN_UP(offset, alignment) \
        (offset) = ((offset) + (alignment) - 1) & ~((alignment) - 1)
    int offset = 0;
    void* ptr;
    ptr = (void*)(size_t)outputDevPtr;
    ALIGN_UP(offset, __alignof(ptr));
    cuParamSetv(cuFunction, offset, &ptr, sizeof(ptr));
    offset += sizeof(ptr);
    ptr = (void*)(size_t)inputDevPtr;
```

```

    ALIGN_UP(offset, __alignof(ptr));
    cuParamSetv(cuFunction, offset, &ptr, sizeof(ptr));
    offset += sizeof(ptr);
    ALIGN_UP(offset, __alignof(size));
    cuParamSeti(cuFunction, offset, size);
    offset += sizeof(int);
    cuParamSetSize(cuFunction, offset);
    cuFuncSetBlockShape(cuFunction, 512, 1, 1);
    cuLaunchGridAsync(cuFunction, 100, 1, stream[i]);
}
for (int i = 0; i < 2; ++i)
    cuMemcpyDtoHAsync(hostPtr + i * size, outputDevPtr + i * size,
                      size, stream[i]);
cuCtxSynchronize();

```

```

for (int i = 0; i < 2; ++i)
    cuStreamDestroy(&stream[i]);

```

3.3.9.2 Event Management

The driver API provides functions similar to the runtime API to manage events. The following code sample is the driver version of the code sample from Section 3.2.6.2.

```

CUevent start, stop;
cuEventCreate(&start);
cuEventCreate(&stop);

```

```

cuEventRecord(start, 0);
for (int i = 0; i < 2; ++i)
    cuMemcpyHtoDAsync(inputDevPtr + i * size, hostPtr + i * size,
                      size, stream[i]);
for (int i = 0; i < 2; ++i) {
    #define ALIGN_UP(offset, alignment) \
        (offset) = ((offset) + (alignment) - 1) & ~((alignment) - 1)
    int offset = 0;
    void* ptr;
    ptr = (void*)(size_t)outputDevPtr;
    ALIGN_UP(offset, __alignof(ptr));
    cuParamSetv(cuFunction, offset, &ptr, sizeof(ptr));
    offset += sizeof(ptr);
    ptr = (void*)(size_t)inputDevPtr;
    ALIGN_UP(offset, __alignof(ptr));
    cuParamSetv(cuFunction, offset, &ptr, sizeof(ptr));
    offset += sizeof(ptr);
    ALIGN_UP(offset, __alignof(size));
    cuParamSeti(cuFunction, offset, size);
    offset += sizeof(size);
    cuParamSetSize(cuFunction, offset);
    cuFuncSetBlockShape(cuFunction, 512, 1, 1);
    cuLaunchGridAsync(cuFunction, 100, 1, stream[i]);
}
for (int i = 0; i < 2; ++i)
    cuMemcpyDtoHAsync(hostPtr + i * size, outputDevPtr + i * size,
                      size, stream[i]);
cuEventRecord(stop, 0);

```

```
cuEventSynchronize(stop);
float elapsedTime;
cuEventElapsedTime(&elapsedTime, start, stop);
```

They are destroyed this way:

```
cuEventDestroy(start);
cuEventDestroy(stop);
```

3.3.9.3 Synchronous Calls

Whether the host thread will yield, block, or spin on a synchronous function call can be specified by calling **cuCtxCreate()** with some specific flags as described in the reference manual.

3.3.10 OpenGL Interoperability

The driver API provides functions similar to the runtime API to manage OpenGL interoperability.

Interoperability with OpenGL requires that the CUDA context be specifically created using **cuGLCtxCreate()** instead of **cuCtxCreate()**.

A buffer object must be registered to CUDA before it can be mapped. This is done with **cuGLRegisterBufferObject()**:

```
GLuint bufferObj;
cuGLRegisterBufferObject(bufferObj);
```

Once it is registered, a buffer object can be read from or written to by kernels using the device memory address returned by **cuGLMapBufferObject()**:

```
GLuint bufferObj;
float* devPtr;
cuGLMapBufferObject((void*)&devPtr, bufferObj);
```

Unmapping is done with **cuGLUnmapBufferObject()** and unregistering with **cuGLUnregisterBufferObject()**.

The following code sample is the driver version of the code sample from Section 3.2.6.3.

```
CUfunction createVertices;
int main()
{
    // Initialize driver API
    ...

    // Get handle for device 0
    CUdevice cuDevice = 0;
    cuDeviceGet(&cuDevice, 0);

    // Create context
    CUcontext cuContext;
    cuGLCtxCreate(&cuContext, 0, cuDevice);

    // Create module from binary file
    CUmodule cuModule;
    cuModuleLoad(&cuModule, "createVertices.ptx");

    // Get function handle from module
```

```

    cuModuleGetFunction(&createVertices,
                       cuModule, "createVertices");

    // Initialize OpenGL and GLUT
    ...
    glutDisplayFunc(display);

    // Create buffer object and register it with CUDA
    glGenBuffers(1, positionsVBO);
    glBindBuffer(GL_ARRAY_BUFFER, &vbo);
    unsigned int size = width * height * 4 * sizeof(float);
    glBufferData(GL_ARRAY_BUFFER, size, 0, GL_DYNAMIC_DRAW);
    glBindBuffer(GL_ARRAY_BUFFER, 0);
    cuGLRegisterBufferObject(positionsVBO);

    // Launch rendering loop
    glutMainLoop();
}

void display()
{
    // Map OpenGL buffer object for writing from CUDA
    CUdeviceptr positions;
    unsigned int size;
    cuGLMapBufferObject(&positions, &size, positionsVBO);

    // Execute kernel
    #define ALIGN_UP(offset, alignment) \
        (offset) = ((offset) + (alignment) - 1) & ~((alignment) - 1)
    int offset = 0;
    void* ptr = (void*)(size_t)positions;
    ALIGN_UP(offset, __alignof(ptr));
    cuParamSetv(createVertices, offset, &ptr, sizeof(ptr));
    offset += sizeof(ptr);
    ALIGN_UP(offset, __alignof(time));
    cuParamSetf(createVertices, offset, time);
    offset += sizeof(time);
    ALIGN_UP(offset, __alignof(width));
    cuParamSeti(createVertices, offset, width);
    offset += sizeof(width);
    ALIGN_UP(offset, __alignof(height));
    cuParamSeti(createVertices, offset, height);
    offset += sizeof(height);
    cuParamSetSize(createVertices, offset);
    int threadsPerBlock = 16;
    cuFuncSetBlockShape(createVertices,
                        threadsPerBlock, threadsPerBlock, 1);
    cuLaunchGrid(createVertices,
                  width / threadsPerBlock, height / threadsPerBlock);

    // Unmap buffer object
    cuGLUnmapBufferObject(positionsVBO);

    // Render from buffer object
    glClear(GL_COLOR_BUFFER_BIT | GL_DEPTH_BUFFER_BIT);
    glBindBuffer(GL_ARRAY_BUFFER, positionsVBO);
    glVertexPointer(4, GL_FLOAT, 0, 0);

```

```

    glEnableClientState(GL_VERTEX_ARRAY);
    glDrawArrays(GL_POINTS, 0, width * height);
    glDisableClientState(GL_VERTEX_ARRAY);

    // Swap buffers
    glutSwapBuffers();
    glutPostRedisplay();
}

void deleteVBO()
{
    cuGLUnregisterBufferObject(positionsVBO);
    glDeleteBuffers(1, &positionsVBO);
}

```

On Windows and for Quadro GPUs, **cuWGLGetDevice()** can be used to retrieve the CUDA device associated to the handle returned by **WGL_NV_gpu_affinity()**.

3.3.11 Direct3D Interoperability

The driver API provides functions similar to the runtime API to manage Direct3D interoperability in the same way.

Interoperability with Direct3D requires that the Direct3D device be specified when the CUDA context is created. This is done by creating the CUDA context using **cuD3D9CtxCreate()** (resp. **cuD3D10CtxCreate()**) instead of **cuCtxCreate()**.

Direct3D resources can then be registered to CUDA using **cuD3D9RegisterResource()** (resp. **cuD3D10RegisterResource()**):

```

IDirect3DVertexBuffer9* buffer;
cuD3D9RegisterResource(buffer, CU_D3D9_REGISTER_FLAGS_NONE);
IDirect3DSurface9* surface;
cuD3D9RegisterResource(surface, CU_D3D9_REGISTER_FLAGS_NONE);

```

```

ID3D10Buffer* buffer;
cuD3D10RegisterResource(buffer, CU_D3D10_REGISTER_FLAGS_NONE);
ID3D10Texture2D* tex2D;
cuD3D10RegisterResource(tex2D, CU_D3D10_REGISTER_FLAGS_NONE);

```

cuD3D9RegisterResource() (resp. **cuD3D10RegisterResource()**) is potentially high-overhead and typically called only once per resource. Unregistering is done with **cuD3D9UnregisterVertexBuffer()** (resp. **cuD3D10UnregisterVertexBuffer()**).

Once a resource is registered to CUDA, it can be mapped and unmapped as many times as necessary using **cuD3D9MapResources()** (resp. **cuD3D10MapResources()**) and **cuD3D9UnmapResources()** (resp. **cuD3D10UnmapResources()**).

A mapped resource can be read from or written to by kernels using the device memory address returned by **cuD3D9ResourceGetMappedPointer()** (resp. **cuD3D10ResourceGetMappedPointer()**) and the size and pitch information returned by **cuD3D9ResourceGetMappedSize()** (resp. **cuD3D10ResourceGetMappedSize()**),

cuD3D9ResourceGetMappedPitch() (resp. **cuD3D10ResourceGetMappedPitch()**), and **cuD3D9ResourceGetMappedPitchSlice()** (resp. **cuD3D10ResourceGetMappedPitchSlice()**).

When applicable, a CUDA array can also be obtained from a mapped resource using **cuD3D9ResourceGetMappedArray()** (resp. **cuD3D10ResourceGetMappedArray()**).

Accessing a resource through Direct3D while it is mapped produces undefined results.

The following code sample is the driver version of the host code of the sample from Section 3.2.8.

```
IDirect3D9* D3D;
IDirect3DDevice9* device;
struct CUSTOMVERTEX {
    FLOAT x, y, z;
    DWORD color;
};
IDirect3DVertexBuffer9* positionsVB;

int main()
{
    // Initialize Direct3D
    D3D = Direct3DCreate9(D3D_SDK_VERSION);

    // Get a CUDA-enabled adapter
    unsigned int adapter = 0;
    for (; adapter < g_pD3D->GetAdapterCount(); adapter++) {
        D3DADAPTER_IDENTIFIER9 adapterId;
        g_pD3D->GetAdapterIdentifier(adapter, 0, &adapterId);
        int dev;
        cuD3D9GetDevice(&dev, adapterId.DeviceName);
        if (cudaSuccess == cudaGetLastError())
            break;
    }

    // Create device
    ...
    D3D->CreateDevice(adapter, D3DDEVTYPE_HAL, hWnd,
                    D3DCREATE_HARDWARE_VERTEXPROCESSING,
                    &params, &device);

    // Initialize driver API
    ...

    // Create context
    CUdevice cuDevice;
    CUcontext cuContext;
    cuD3D9CtxCreate(&cuContext, &cuDevice, 0, &device);

    // Create module from binary file
    CUmodule cuModule;
    cuModuleLoad(&cuModule, "createVertices.ptx");
}
```

```

// Get function handle from module
cuModuleGetFunction(&createVertices,
                    cuModule, "createVertices");

// Create vertex buffer and register it with CUDA
unsigned int size = width * height * sizeof(CUSTOMVERTEX);
device->CreateVertexBuffer(size, 0, D3DFVF_CUSTOMVERTEX,
                           D3DP00L_DEFAULT, &positionsVB, 0);
cuD3D9RegisterResource(positionsVB,
                        CU_D3D9_REGISTER_FLAGS_NONE);
cuD3D9ResourceSetMapFlags(positionsVB,
                           CU_D3D9_MAPRESOURCE_FLAGS_WRITEDISCARD);

// Launch rendering loop
while (...) {
    ...
    Render();
    ...
}

void Render()
{
    // Map vertex buffer for writing from CUDA
    float4* positions;
    cuD3D9MapResources(1, (IDirect3DResource9**)&positionsVB);
    cuD3D9ResourceGetMappedPointer((void**)&positions,
                                    positionsVB, 0, 0);

    // Execute kernel
    #define ALIGN_UP(offset, alignment) \
        (offset) = ((offset) + (alignment) - 1) & ~((alignment) - 1)
    int offset = 0;
    void* ptr = (void*)(size_t)positions;
    ALIGN_UP(offset, __alignof(ptr));
    cuParamSetv(createVertices, offset, &ptr, sizeof(ptr));
    offset += sizeof(ptr);
    ALIGN_UP(offset, __alignof(time));
    cuParamSetf(createVertices, offset, time);
    offset += sizeof(time);
    ALIGN_UP(offset, __alignof(width));
    cuParamSeti(createVertices, offset, width);
    offset += sizeof(width);
    ALIGN_UP(offset, __alignof(height));
    cuParamSeti(createVertices, offset, height);
    offset += sizeof(height);
    cuParamSetSize(createVertices, offset);
    int threadsPerBlock = 16;
    cuFuncSetBlockShape(createVertices,
                        threadsPerBlock, threadsPerBlock, 1);
    cuLaunchGrid(createVertices,
                  width / threadsPerBlock, height / threadsPerBlock);

    // Unmap vertex buffer
    cuD3D9UnmapResources(1, (IDirect3DResource9**)&positionsVB);

    // Draw and present

```

```

    ...
}

void releaseVB()
{
    cuD3D9UnregisterResource(positionsVB);
    positionsVB->Release();
}

ID3D10Device* device;
struct CUSTOMVERTEX {
    FLOAT x, y, z;
    DWORD color;
};
ID3D10Buffer* positionsVB;

int main()
{
    // Get a CUDA-enabled adapter
    IDXGIFactory* factory;
    CreatedXGIFactory(__uuidof(IDXGIFactory), (void**)&factory);
    IDXGIAdapter* adapter = 0;
    for (unsigned int i = 0; !adapter; ++i) {
        if (FAILED(factory->EnumAdapters(i, &adapter))
            break;
        int dev;
        cuD3D10GetDevice(&dev, adapter);
        if (cudaSuccess == cudaGetLastError())
            break;
        adapter->Release();
    }
    factory->Release();

    // Create swap chain and device
    ...
    D3D10CreateDeviceAndSwapChain(adapter,
                                   D3D10_DRIVER_TYPE_HARDWARE, 0,
                                   D3D10_CREATE_DEVICE_DEBUG,
                                   D3D10_SDK_VERSION,
                                   &swapChainDesc &swapChain,
                                   &device);

    adapter->Release();

    // Initialize driver API
    ...

    // Create context
    CUdevice cuDevice;
    CUcontext cuContext;
    cuD3D10CtxCreate(&cuContext, &cuDevice, 0, &device);

    // Create module from binary file
    CUmodule cuModule;
    cuModuleLoad(&cuModule, "createVertices.ptx");

    // Get function handle from module

```

```

cuModuleGetFunction(&createVertices,
                    cuModule, "createVertices");

// Create vertex buffer and register it with CUDA
unsigned int size = width * height * sizeof(CUSTOMVERTEX);
D3D10_BUFFER_DESC bufferDesc;
bufferDesc.Usage          = D3D10_USAGE_DEFAULT;
bufferDesc.ByteWidth      = size;
bufferDesc.BindFlags      = D3D10_BIND_VERTEX_BUFFER;
bufferDesc.CPUAccessFlags = 0;
bufferDesc.MiscFlags      = 0;
device->CreateBuffer(&bufferDesc, 0, &positionsVB);
cuD3D10RegisterResource(positionsVB,
                        CU_D3D10_REGISTER_FLAGS_NONE);
cuD3D10ResourceSetMapFlags(positionsVB,
                           CU_D3D10_MAPRESOURCE_FLAGS_WRITEDISCARD);

// Launch rendering loop
while (...) {
    ...
    Render();
    ...
}
}

void Render()
{
    // Map vertex buffer for writing from CUDA
    float4* positions;
    cuD3D10MapResources(1, (ID3D10Resource**)&positionsVB);
    cuD3D10ResourceGetMappedPointer((void**)&positions,
                                    positionsVB, 0);

    // Execute kernel
    #define ALIGN_UP(offset, alignment) \
        (offset) = ((offset) + (alignment) - 1) & ~((alignment) - 1)
    int offset = 0;
    void* ptr = (void*)(size_t)positions;
    ALIGN_UP(offset, __alignof(ptr));
    cuParamSetv(createVertices, offset, &ptr, sizeof(ptr));
    offset += sizeof(ptr);
    ALIGN_UP(offset, __alignof(time));
    cuParamSetf(createVertices, offset, time);
    offset += sizeof(time);
    ALIGN_UP(offset, __alignof(width));
    cuParamSeti(createVertices, offset, width);
    offset += sizeof(width);
    ALIGN_UP(offset, __alignof(height));
    cuParamSeti(createVertices, offset, height);
    offset += sizeof(height);
    cuParamSetSize(createVertices, offset);
    int threadsPerBlock = 16;
    cuFuncSetBlockShape(createVertices,
                        threadsPerBlock, threadsPerBlock, 1);
    cuLaunchGrid(createVertices,
                 width / threadsPerBlock, height / threadsPerBlock);
}

```

```

    // Unmap vertex buffer
    cuD3D10UnmapResources(1, (ID3D10Resource**)&positionsVB);

    // Draw and present
    ...
}

void releaseVB()
{
    cuD3D10UnregisterResource(positionsVB);
    positionsVB->Release();
}

```

In the following code sample, each thread accesses one pixel of a 2D surface of size **(width, height)** and pixel format **float4**:

```

// host code
CUdeviceptr devPtr;
cuD3D9ResourceGetMappedPointer(&devPtr, surface, 0, 0);
size_t pitch;
cuD3D9ResourceGetMappedPitch(&pitch, 0, surface, 0, 0);
cuModuleGetFunction(&cuFunction, cuModule, "myKernel");
#define ALIGN_UP(offset, alignment) \
    ((offset) + (alignment) - 1) & ~((alignment) - 1)
int offset = 0;
void* ptr = (void*)(size_t)devPtr;
ALIGN_UP(offset, __alignof(ptr));
cuParamSetv(cuFunction, offset, &ptr, sizeof(ptr));
offset += sizeof(ptr);
ALIGN_UP(offset, __alignof(width));
cuParamSeti(cuFunction, offset, width);
offset += sizeof(width);
ALIGN_UP(offset, __alignof(height));
cuParamSeti(cuFunction, offset, height);
offset += sizeof(height);
ALIGN_UP(offset, __alignof(pitch));
cuParamSeti(cuFunction, offset, pitch);
offset += sizeof(pitch);
cuParamSetSize(cuFunction, offset);
cuFuncSetBlockShape(cuFunction, 16, 16, 1);
cuLaunchGrid(cuFunction,
              (width+Db.x-1)/Db.x, (height+Db.y-1)/Db.y);

// device code
__global__ void myKernel(unsigned char* surface,
                        int width, int height, size_t pitch)
{
    int x = blockIdx.x * blockDim.x + threadIdx.x;
    int y = blockIdx.y * blockDim.y + threadIdx.y;
    if (x >= width || y >= height) return;
    float* pixel = (float*)(surface + y * pitch) + 4 * x;
}

// host code
CUdeviceptr devPtr;
cuD3D10ResourceGetMappedPointer(&devPtr, surface, 0);
size_t pitch;
cuD3D10ResourceGetMappedPitch(&pitch, 0, surface, 0);

```

```

cuModuleGetFunction(&cuFunction, cuModule, "myKernel");
#define ALIGN_UP(offset, alignment) \
    (offset) = ((offset) + (alignment) - 1) & ~((alignment) - 1)
int offset = 0;
void* ptr = (void*)(size_t)devPtr;
ALIGN_UP(offset, __alignof(ptr));
cuParamSetv(cuFunction, offset, &ptr, sizeof(ptr));
offset += sizeof(ptr);
ALIGN_UP(offset, __alignof(width));
cuParamSeti(cuFunction, offset, width);
offset += sizeof(width);
ALIGN_UP(offset, __alignof(height));
cuParamSeti(cuFunction, offset, height);
offset += sizeof(height);
ALIGN_UP(offset, __alignof(pitch));
cuParamSeti(cuFunction, offset, pitch);
offset += sizeof(pitch);
cuParamSetSize(cuFunction, offset);
cuFuncSetBlockShape(cuFunction, 16, 16, 1);
cuLaunchGrid(cuFunction,
              (width+Db.x-1)/Db.x, (height+Db.y-1)/Db.y);

// device code
__global__ void myKernel(unsigned char* surface,
                          int width, int height, size_t pitch)
{
    int x = blockIdx.x * blockDim.x + threadIdx.x;
    int y = blockIdx.y * blockDim.y + threadIdx.y;
    if (x >= width || y >= height) return;
    float* pixel = (float*)(surface + y * pitch) + 4 * x;
}

```

3.3.12 Error Handling

All driver functions return an error code, but for an asynchronous function (see Section 3.2.6), this error code cannot possibly report any of the asynchronous errors that could occur on the device since the function returns before the device has completed the task; the error code only reports errors that occur on the host prior to executing the task, typically related to parameter validation; if an asynchronous error occurs, it will be reported by some subsequent unrelated runtime function call.

The only way to check for asynchronous errors just after some asynchronous function call is therefore to synchronize just after the call by calling **cuCtxSynchronize()** (or by using any other synchronization mechanisms described in Section 3.3.9) and checking the error code returned by **cuCtxSynchronize()**.

3.4 Versioning and Compatibility

There are two version numbers that developers should care about when developing a CUDA application: The compute capability that describes the general specifications and features of the compute device (see Section 2.5) and the version

of the CUDA driver API that describes the features supported by the driver API and runtime.

The version of the driver API is defined in the driver header file as **CUDA_VERSION**. It allows developers to check whether their application requires a newer driver than the one currently installed. This is important, because the driver API is *backward compatible*, meaning that applications, plug-ins, and libraries (including the C runtime) compiled against a particular version of the driver API will continue to work on subsequent driver releases as illustrated in Figure 3-4. The driver API is not *forward compatible*, which means that applications, plug-ins, and libraries (including the C runtime) compiled against a particular version of the driver API will not work on previous versions of the driver.

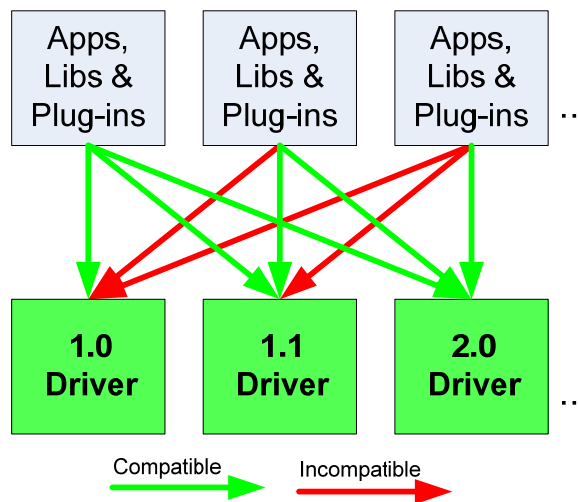


Figure 3-4. The Driver API is Backward, but Not Forward Compatible

It is important to note that mixing and matching versions is not supported; specifically:

- ❑ All applications, plug-ins, and libraries on a system must use the same version of the CUDA driver API, since only one version of the CUDA driver can be installed on a system.
- ❑ All plug-ins and libraries used by an application must use the same version of the runtime.
- ❑ All plug-ins and libraries used by an application must use the same version of any libraries that use the runtime (such as CUFFT, CUBLAS, ...).

3.5 Compute Modes

On Tesla solutions running Linux, one can set any device in a system in one of the three following modes using NVIDIA's System Management Interface (`nvidia-smi`), which is a tool distributed as part of the Linux driver:

- ❑ *Default* compute mode: Multiple host threads can use the device (by calling **cudaSetDevice()** on this device, when using the runtime API, or by making current a context associated to the device, when using the driver API) at the same time.
- ❑ *Exclusive* compute mode: Only one host thread can use the device at any given time.
- ❑ *Prohibited* compute mode: No host thread can use the device.

This means, in particular, that a host thread using the runtime API without explicitly calling **cudaSetDevice()** might be associated with a device other than device 0 if device 0 turns out to be in prohibited compute mode or in exclusive compute mode and used by another host thread. **cudaSetValidDevices()** can be used to set a device from a prioritized list of devices.

Applications may query the compute mode of a device by calling **cudaGetDeviceProperties()** and checking the **computeMode** property or checking the **CU_DEVICE_COMPUTE_MODE** attribute using **cuDeviceGetAttribute()**.

3.6 Mode Switches

GPUs dedicate some DRAM memory to the so-called *primary surface*, which is used to refresh the display device whose output is viewed by the user. When users initiate a *mode switch* of the display by changing the resolution or bit depth of the display (using NVIDIA control panel or the Display control panel on Windows), the amount of memory needed for the primary surface changes. For example, if the user changes the display resolution from 1280x1024x32-bit to 1600x1200x32-bit, the system must dedicate 7.68 MB to the primary surface rather than 5.24 MB. (Full-screen graphics applications running with anti-aliasing enabled may require much more display memory for the primary surface.) On Windows, other events that may initiate display mode switches include launching a full-screen DirectX application, hitting Alt+Tab to task switch away from a full-screen DirectX application, or hitting Ctrl+Alt+Del to lock the computer.

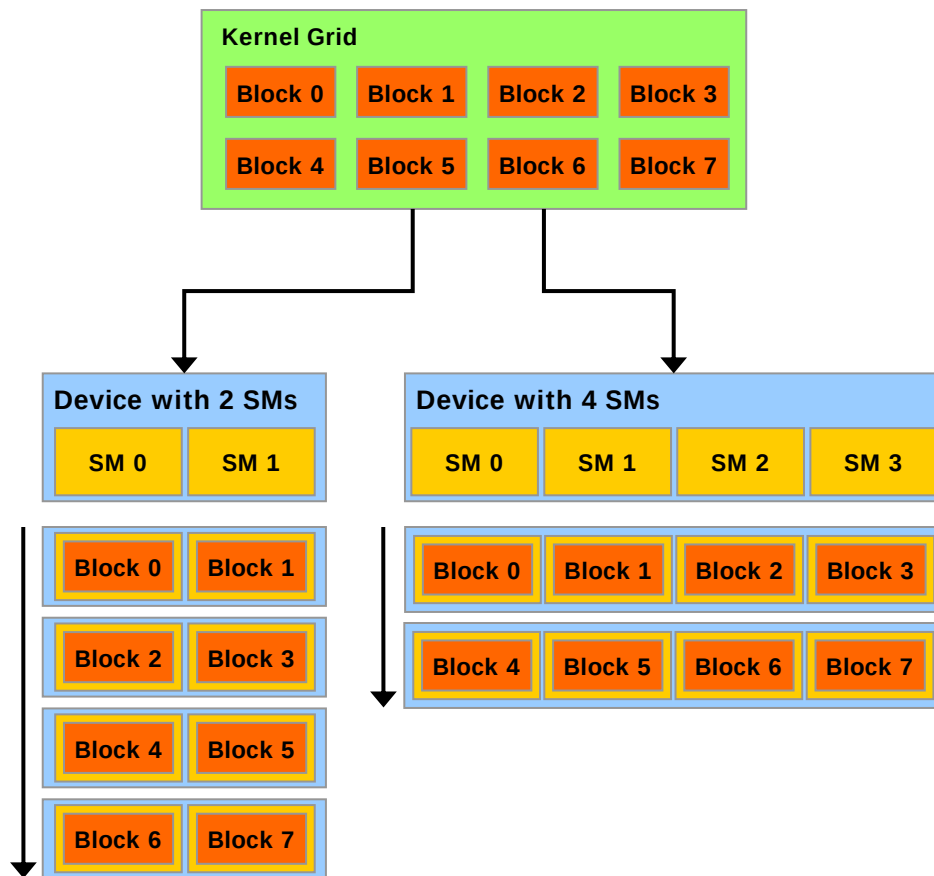
If a mode switch increases the amount of memory needed for the primary surface, the system may have to cannibalize memory allocations dedicated to CUDA applications. Therefore, a mode switch results in any call to the CUDA runtime to fail and return an invalid context error.



Chapter 4. Hardware Implementation

4.1 A Set of SIMT Multiprocessors with On-Chip Shared Memory

The CUDA architecture is built around a scalable array of multithreaded Streaming Multiprocessors (SMs). When a CUDA program on the host CPU invokes a kernel grid, the blocks of the grid are enumerated and distributed to multiprocessors with available execution capacity as illustrated in Figure 4-1. The threads of a thread block execute concurrently on one multiprocessor. As thread blocks terminate, new blocks are launched on the vacated multiprocessors.



A device with more multiprocessors will automatically execute a kernel grid in less time than a device with fewer multiprocessors.

Figure 4-1. Automatic Scalability

A multiprocessor consists of eight Scalar Processor (SP) cores, two special function units for transcendentals, a multithreaded instruction unit, and on-chip shared memory. The multiprocessor creates, manages, and executes concurrent threads in hardware with zero scheduling overhead. It implements the `__syncthreads()` barrier synchronization intrinsic with a single instruction. Fast barrier synchronization together with lightweight thread creation and zero-overhead thread scheduling efficiently support very fine-grained parallelism, allowing, for example, a low granularity decomposition of problems by assigning one thread to each data element (such as a pixel in an image, a voxel in a volume, a cell in a grid-based computation).

To manage hundreds of threads running several different programs, the multiprocessor employs a new architecture we call SIMT (single-instruction, multiple-thread). The multiprocessor maps each thread to one scalar processor core, and each scalar thread executes independently with its own instruction address and register state. The multiprocessor SIMT unit creates, manages, schedules, and executes threads in groups of 32 parallel threads called *warps*. (This term originates from weaving, the first parallel thread technology. A *half-warp* is either the first or second half of a warp.) Individual threads composing a SIMT warp start together at

the same program address but are otherwise free to branch and execute independently.

When a multiprocessor is given one or more thread blocks to execute, it splits them into warps that get scheduled by the SIMT unit. The way a block is split into warps is always the same; each warp contains threads of consecutive, increasing thread IDs with the first warp containing thread 0. Section 2.2 describes how thread IDs relate to thread indices in the block.

Every instruction issue time, the SIMT unit selects a warp that is ready to execute and issues the next instruction to the active threads of the warp. A warp executes one common instruction at a time, so full efficiency is realized when all 32 threads of a warp agree on their execution path. If threads of a warp diverge via a data-dependent conditional branch, the warp serially executes each branch path taken, disabling threads that are not on that path, and when all paths complete, the threads converge back to the same execution path. Branch divergence occurs only within a warp; different warps execute independently regardless of whether they are executing common or disjointed code paths.

SIMT architecture is akin to SIMD (Single Instruction, Multiple Data) vector organizations in that a single instruction controls multiple processing elements. A key difference is that SIMD vector organizations expose the SIMD width to the software, whereas SIMT instructions specify the execution and branching behavior of a single thread. In contrast with SIMD vector machines, SIMT enables programmers to write thread-level parallel code for independent, scalar threads, as well as data-parallel code for coordinated threads. For the purposes of correctness, the programmer can essentially ignore the SIMT behavior; however, substantial performance improvements can be realized by taking care that the code seldom requires threads in a warp to diverge. In practice, this is analogous to the role of cache lines in traditional code: Cache line size can be safely ignored when designing for correctness but must be considered in the code structure when designing for peak performance. Vector architectures, on the other hand, require the software to coalesce loads into vectors and manage divergence manually.

As illustrated by Figure 4-2, each multiprocessor has on-chip memory of the four following types:

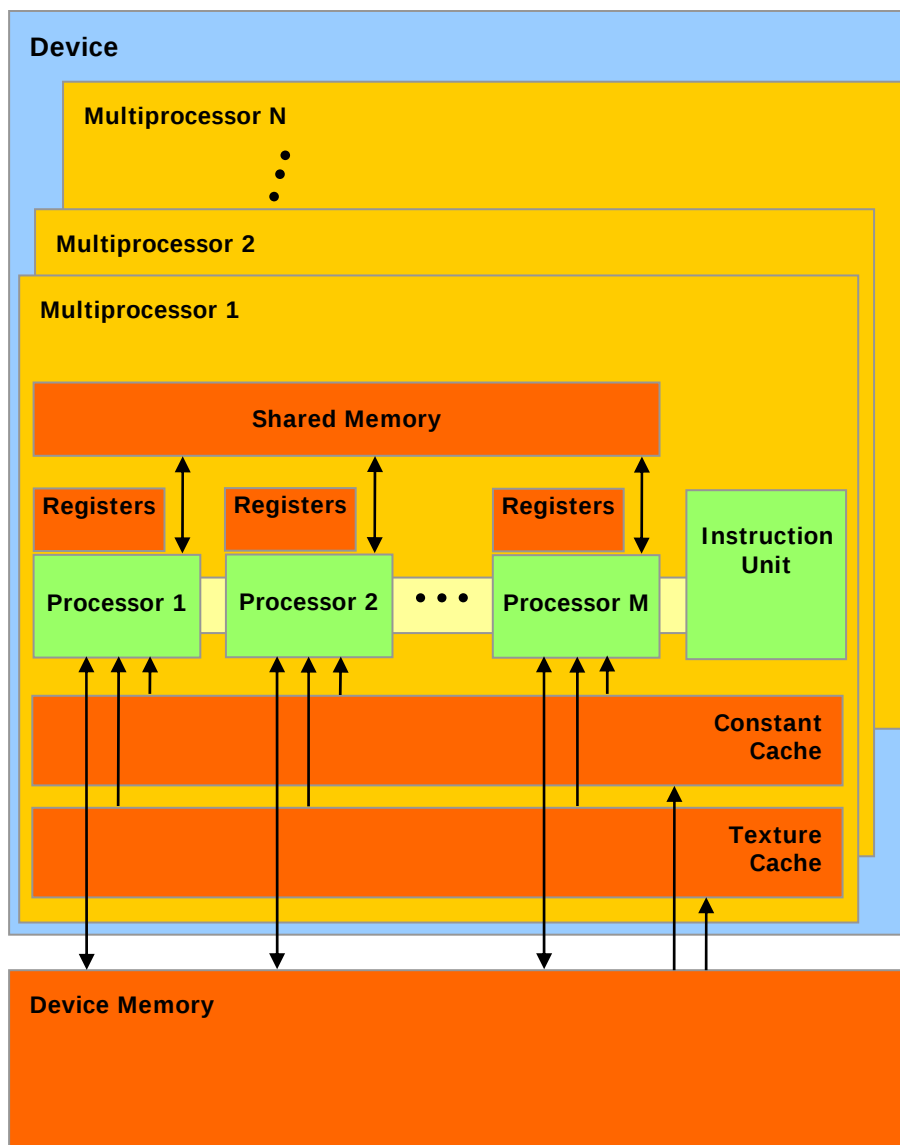
- ❑ One set of local 32-bit *registers* per processor,
- ❑ A parallel data cache or *shared memory* that is shared by all scalar processor cores and is where the shared memory space resides,
- ❑ A read-only *constant cache* that is shared by all scalar processor cores and speeds up reads from the constant memory space, which is a read-only region of device memory,
- ❑ A read-only *texture cache* that is shared by all scalar processor cores and speeds up reads from the texture memory space, which is a read-only region of device memory; each multiprocessor accesses the texture cache via a *texture unit* that implements the various addressing modes and data filtering mentioned in Section 3.2.4.

The local and global memory spaces are read-write regions of device memory and are not cached.

The number of blocks a multiprocessor can process at once – referred to as the number of *active* blocks per multiprocessor – depends on how many registers per

thread and how much shared memory per block are required for a given kernel since the multiprocessor's registers and shared memory are split among all the threads of the active blocks. If there are not enough registers or shared memory available per multiprocessor to process at least one block, the kernel will fail to launch. The maximum number of active blocks per multiprocessor, as well as the maximum number of active warps and maximum number of active threads are given in Appendix A.

If a non-atomic instruction executed by a warp writes to the same location in global or shared memory for more than one of the threads of the warp, the number of serialized writes that occur to that location and the order in which they occur is undefined, but one of the writes is guaranteed to succeed. If an atomic instruction (see Section B.10) executed by a warp reads, modifies, and writes to the same location in global memory for more than one of the threads of the warp, each read, modify, write to that location occurs and they are all serialized, but the order in which they occur is undefined.



A set of SIMT multiprocessors with on-chip shared memory.

Figure 4-2. Hardware Model

4.2 Multiple Devices

The use of multiple GPUs as CUDA devices by an application running on a multi-GPU system is only guaranteed to work if these GPUs are of the same type.

When the system is in SLI mode, all GPUs are accessible via the CUDA driver and runtime as separate devices, but there are special considerations as described below.

First, an allocation in one CUDA device on one GPU will consume memory on other GPUs. Because of this, allocations may fail earlier than otherwise expected.

Second, when a Direct3D application runs in SLI Alternate Frame Rendering mode, the Direct3D device(s) created by that application can be used for CUDA-Direct3D interoperability (i.e. passed as a parameter to

cudaD3D[9|10]SetDirect3DDevice() when using the runtime API), but only one CUDA device can be created at a time from one of these Direct3D devices.

This CUDA device only executes the CUDA work on one of the GPUs in the SLI configuration. As a consequence, real interoperability only happens with the copy of a Direct3D resource in that GPU (Note: in AFR mode Direct3D resources that must be in GPU memory are duplicated in the GPU memory of each GPU in the SLI configuration). In some cases this is not the desired behavior and an application may need to forfeit CUDA-Direct3D interoperability and manually copy the output of its CUDA work to Direct3D resources using the existing CUDA and Direct3D API.

Chapter 5.

Performance Guidelines

5.1 Instruction Performance

To process an instruction for a warp of threads, a multiprocessor must:

- ❑ Read the instruction operands for each thread of the warp,
- ❑ Execute the instruction,
- ❑ Write the result for each thread of the warp.

Therefore, the effective instruction throughput depends on the nominal instruction throughput as well as the memory latency and bandwidth. It is maximized by:

- ❑ Minimizing the use of instructions with low throughput (see Section 5.1.1),
- ❑ Maximizing the use of the available memory bandwidth for each category of memory (see Section 5.1.2),
- ❑ Allowing the thread scheduler to overlap memory transactions with mathematical computations as much as possible, which requires that:
 - The program executed by the threads is of high arithmetic intensity, that is, has a high number of arithmetic operations per memory operation;
 - There are many active threads per multiprocessor, as detailed in Section 5.2.

5.1.1 Instruction Throughput

In this section, throughputs are given in number of operations per clock cycle per multiprocessor. For a warp size of 32, an instruction is made of 32 operations. Therefore, if T is the number of operations per clock cycle, the instruction throughput is one instruction every $32/T$ clock cycles.

All throughputs are for one multiprocessor. They must be multiplied by the number of multiprocessors in the device to get throughput for the whole device.

5.1.1.1 Arithmetic Instructions

For single-precision floating-point code, we highly recommend use of the **float** type and the single-precision floating-point mathematical functions. When compiling for devices without native double-precision floating-point support, such as devices of compute capability 1.2 and lower, each **double** variable gets converted to single-precision floating-point format (but retains its size of 64 bits)

and double-precision floating-point arithmetic gets demoted to single-precision floating-point arithmetic.

Single-Precision Floating-Point Basic Arithmetic

Throughput of single-precision floating-point add, multiply, and multiply-add is 8 operations per clock cycle.

Throughput of reciprocal is 2 operations per clock cycle.

Throughput of single-precision floating-point division is 0.88 operations per clock cycle, but `__fdivdef(x, y)` (see Section C.2) provides a faster version with a throughput of 1.6 operations per clock cycle.

Single-Precision Floating-Point Square Root and Reciprocal Square Root

Throughput of reciprocal square root is 2 operations per clock cycle.

Single-precision floating-point square root is implemented as a reciprocal square root followed by a reciprocal instead of a reciprocal square root followed by a multiplication, so that it gives correct results for 0 and infinity. Therefore, its throughput is 1 operation per clock cycle.

Single-Precision Floating-Point Logarithm

Throughput of `__logf(x)` (see Section C.2) is 2 operations per clock cycle.

Sine and Cosine

Throughput of `__sinf(x)`, `__cosf(x)`, `__expf(x)` (see Section C.2) is 1 operation per clock cycle.

`sinf(x)`, `cosf(x)`, `tanf(x)`, `sincosf(x)` and corresponding double-precision instructions are much more expensive and even more so if the absolute value of **`x`** needs to be reduced.

More precisely, the argument reduction code (see `math_functions.h` for implementation) comprises two code paths referred to as the fast path and the slow path, respectively.

The fast path is used for arguments sufficiently small in magnitude and essentially consists of a few multiply-add operations. The slow path is used for arguments large in magnitude, and consists of lengthy computations required to achieve correct results over the entire argument range.

At present, the argument reduction code for the trigonometric functions selects the fast path for arguments whose magnitude is less than 48039.0f for the single-precision functions, and less than 2147483648.0 for the double-precision functions.

As the slow path requires more registers than the fast path, an attempt has been made to reduce register pressure in the slow path by storing some intermediate variables in local memory, which may affect performance because of local memory high latency and bandwidth (see Section 5.1.2.2). At present, 28 bytes of local memory are used by single-precision functions, and 44 bytes are used by double-precision functions. However, the exact amount is subject to change.

Due to the lengthy computations and use of local memory in the slow path, the trigonometric functions throughput is lower by one order of magnitude when the slow path reduction is used as opposed to the fast path reduction.

Integer Arithmetic

Throughput of integer add is 8 operations per clock cycle.

Throughput of 32-bit integer multiplication is 2 operations per clock cycle, but `__mul124` and `__umul124` (see Section C.2) provide signed and unsigned 24-bit integer multiplication with a throughput of 8 operations per clock cycle. On future architectures however, `__[u]mul124` will be slower than 32-bit integer multiplication, so we recommend to provide two kernels, one using `__[u]mul124` and the other using generic 32-bit integer multiplication, to be called appropriately by the application.

Integer division and modulo operation are particularly costly and should be avoided if possible or replaced with bitwise operations whenever possible: If `n` is a power of 2, `(i/n)` is equivalent to `(i>>log2(n))` and `(i%n)` is equivalent to `(i&(n-1))`; the compiler will perform these conversions if `n` is literal.

Comparison

Throughput of compare, min, max is 8 operations per clock cycle.

Bitwise Operations

Throughput of any bitwise operation is 8 operations per clock cycle.

Type Conversion

Throughput of type conversion operations is 8 operations per clock cycle.

Sometimes, the compiler must insert conversion instructions, introducing additional execution cycles. This is the case for:

- ❑ Functions operating on **char** or **short** whose operands generally need to be converted to **int**,
- ❑ Double-precision floating-point constants (defined without any type suffix) used as input to single-precision floating-point computations.

This last case can be avoided by using single-precision floating-point constants, defined with an **f** suffix such as `3.141592653589793f`, `1.0f`, `0.5f`.

5.1.1.2 Control Flow Instructions

Any flow control instruction (**if**, **switch**, **do**, **for**, **while**) can significantly impact the effective instruction throughput by causing threads of the same warp to diverge, that is, to follow different execution paths. If this happens, the different executions paths have to be serialized, increasing the total number of instructions executed for this warp. When all the different execution paths have completed, the threads converge back to the same execution path.

To obtain best performance in cases where the control flow depends on the thread ID, the controlling condition should be written so as to minimize the number of divergent warps. This is possible because the distribution of the warps across the block is deterministic as mentioned in Section 4.1. A trivial example is when the controlling condition only depends on `(threadIdx / warpSize)` where

warpSize is the warp size. In this case, no warp diverges since the controlling condition is perfectly aligned with the warps.

Sometimes, the compiler may unroll loops or it may optimize out **if** or **switch** statements by using branch predication instead, as detailed below. In these cases, no warp can ever diverge. The programmer can also control loop unrolling using the **#pragma unroll** directive (see Section 3.1.2).

When using branch predication none of the instructions whose execution depends on the controlling condition gets skipped. Instead, each of them is associated with a per-thread condition code or *predicate* that is set to true or false based on the controlling condition and although each of these instructions gets scheduled for execution, only the instructions with a true predicate are actually executed. Instructions with a false predicate do not write results, and also do not evaluate addresses or read operands.

The compiler replaces a branch instruction with predicated instructions only if the number of instructions controlled by the branch condition is less or equal to a certain threshold: If the compiler determines that the condition is likely to produce many divergent warps, this threshold is 7, otherwise it is 4.

5.1.1.3 Memory Instructions

Memory instructions include any instruction that reads from or writes to shared, local or global memory. Local memory accesses only occur for some automatic variables as detailed in Section B.2.4.

Throughput of memory operations is 8 operations per clock cycle. When accessing local or global memory, there are, in addition, 400 to 600 clock cycles of memory latency.

As an example, the throughput for the assignment operator in the following sample code:

```
__shared__ float shared[32];
__device__ float device[32];
shared[threadIdx.x] = device[threadIdx.x];
```

is 8 operations per clock cycle for the read from global memory, 8 operations per clock cycle for the write to shared memory, but above all, there is a latency of 400 to 600 clock cycles to read data from global memory.

Much of this global memory latency can be hidden by the thread scheduler if there are sufficient independent arithmetic instructions that can be issued while waiting for the global memory access to complete.

5.1.1.4 Synchronization Instruction

Throughput for **__syncthreads** is 8 operations per clock cycle in the case where no thread has to wait for any other threads.

5.1.2 Memory Bandwidth

The effective bandwidth of each memory space depends significantly on the memory access pattern as detailed in the following sub-sections.

Since device memory is of much higher latency and lower bandwidth than on-chip memory, device memory accesses should be minimized. A typical programming

pattern is to stage data coming from device memory into shared memory; in other words, to have each thread of a block:

- ❑ Load data from device memory to shared memory,
- ❑ Synchronize with all the other threads of the block so that each thread can safely read shared memory locations that were written by different threads,
- ❑ Process the data in shared memory,
- ❑ Synchronize again if necessary to make sure that shared memory has been updated with the results,
- ❑ Write the results back to device memory.

5.1.2.1 Global Memory

The global memory space is not cached, so it is all the more important to follow the right access pattern to get maximum memory bandwidth, especially given how costly accesses to device memory are.

First, the device is capable of reading 4-byte, 8-byte, or 16-byte words from global memory into registers in a single instruction. To have assignments such as:

```
__device__ type device[32];
type data = device[tid];
```

compile to a single load instruction, **type** must be such that **sizeof(type)** is equal to 4, 8, or 16 and variables of type **type** must be aligned to **sizeof(type)** bytes (that is, have their address be a multiple of **sizeof(type)**).

The alignment requirement is automatically fulfilled for built-in types of Section B.3.1 like **float2** or **float4**.

For structures, the size and alignment requirements can be enforced by the compiler using the alignment specifiers **__align__(8)** or **__align__(16)**, such as

```
struct __align__(8) {
    float a;
    float b;
};
```

or

```
struct __align__(16) {
    float a;
    float b;
    float c;
};
```

For structures larger than 16 bytes, the compiler generates several load instructions. To ensure that it generates the minimum number of instructions, such structures should be defined with **__align__(16)**, such as

```
struct __align__(16) {
    float a;
    float b;
    float c;
    float d;
    float e;
};
```

which is compiled into two 16-byte load instructions instead of five 8-byte load instructions.

Any address of a variable residing in global memory or returned by one of the memory allocation routines from the driver or runtime API is always aligned to at least 256 bytes.

Reading mis-aligned 8-byte or 16-byte words produces incorrect results (off by a few words), so special care must be taken to maintain alignment of the starting address of any value or array of values of these types. A typical case where this might be easily overlooked is when using some custom global memory allocation scheme, whereby the allocations of multiple arrays (with multiple calls to **cudaMalloc()** or **cuMemAlloc()**) is replaced by the allocation of a single large block of memory partitioned into multiple arrays, in which case the starting address of each array is offset from the block's starting address.

Second, global memory bandwidth is used most efficiently when the simultaneous memory accesses by threads in a half-warp (during the execution of a single read or write instruction) can be *coalesced* into a single memory transaction of 32, 64, or 128 bytes.

The rest of this section describes the various requirements for memory accesses to coalesce based on the compute capability of the device. If a half-warp fulfills these requirements, coalescing is achieved even if the warp is divergent and some threads of the half-warp do not actually access memory.

For the purpose of the following discussion, global memory is considered to be partitioned into segments of size equal to 32, 64, or 128 bytes *and* aligned to this size.

Coalescing on Devices with Compute Capability 1.0 and 1.1

The global memory access by all threads of a half-warp is coalesced into one or two memory transactions if it satisfies the following three conditions:

- ❑ Threads must access
 - ❑ Either 4-byte words, resulting in one 64-byte memory transaction,
 - ❑ Or 8-byte words, resulting in one 128-byte memory transaction,
 - ❑ Or 16-byte words, resulting in two 128-byte memory transactions;
- ❑ All 16 words must lie in the same segment of size equal to the memory transaction size (or twice the memory transaction size when accessing 16-byte words);
- ❑ Threads must access the words in sequence: The k^{th} thread in the half-warp must access the k^{th} word.

If a half-warp does not fulfill all the requirements above, a separate memory transaction is issued for each thread and throughput is significantly reduced.

Figure 5-1 shows some examples of coalesced memory accesses, while Figure 5-2 and Figure 5-3 show some examples of memory accesses that are non-coalesced for devices of compute capability 1.0 or 1.1.

Coalesced 8-byte accesses deliver a little lower bandwidth than coalesced 4-byte accesses and coalesced 16-byte accesses deliver a noticeably lower bandwidth than coalesced 4-byte accesses. But, while bandwidth for non-coalesced accesses is around an order of magnitude lower than for coalesced accesses when these accesses are 4-byte, it is only around four times lower when they are 8-byte and around two times when they are 16-byte.

Coalescing on Devices with Compute Capability 1.2 and Higher

The global memory access by all threads of a half-warp is coalesced into a single memory transaction as soon as the words accessed by all threads lie in the same segment of size equal to:

- ❑ 32 bytes if all threads access 1-byte words,
- ❑ 64 bytes if all threads access 2-byte words,
- ❑ 128 bytes if all threads access 4-byte or 8-byte words.

Coalescing is achieved for any pattern of addresses requested by the half-warp, including patterns where multiple threads access the same address. This is in contrast with devices of lower compute capabilities where threads need to access words in sequence.

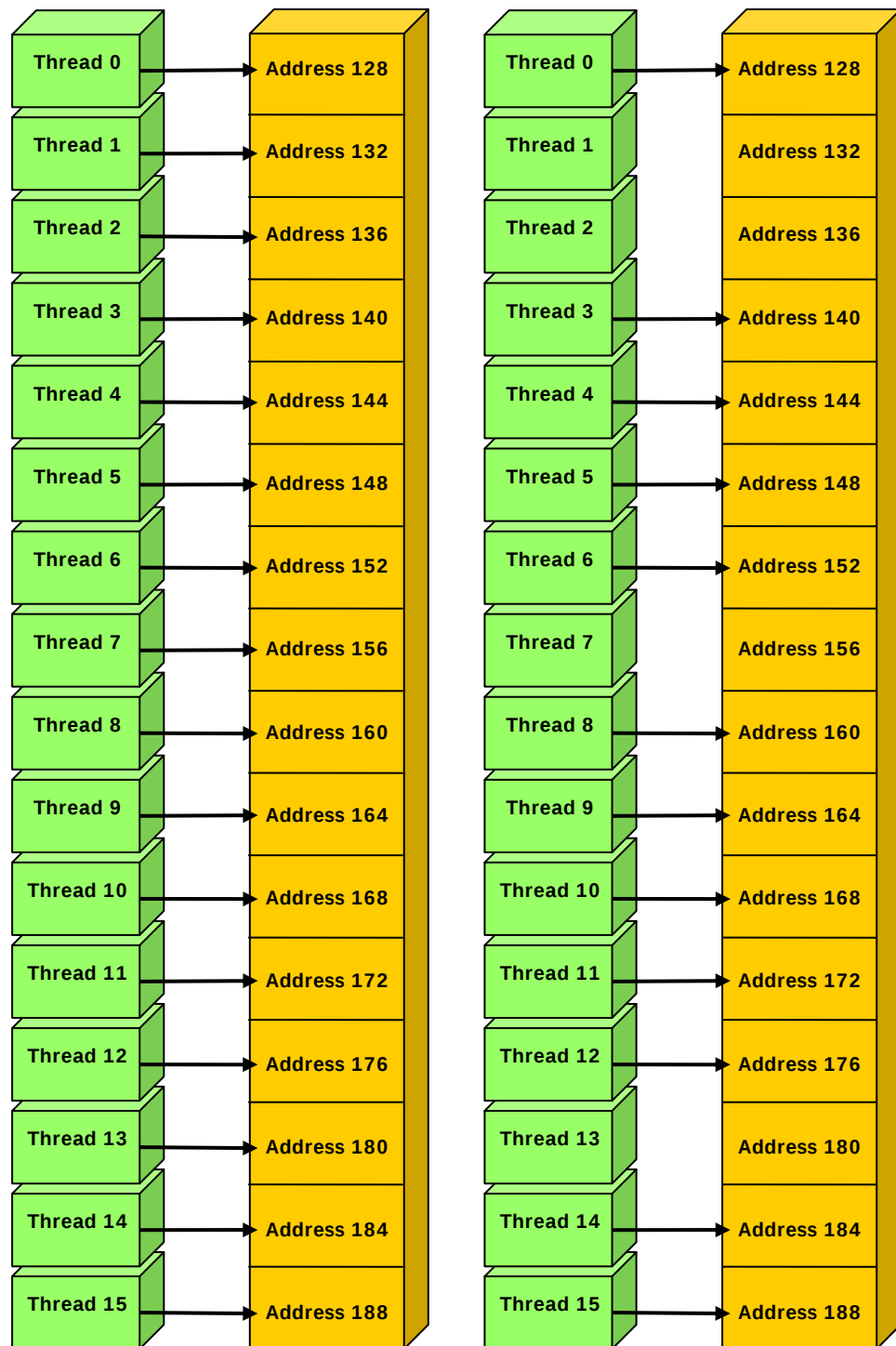
If a half-warp addresses words in n different segments, n memory transactions are issued (one for each segment), whereas devices with lower compute capabilities would issue 16 transactions as soon as n is greater than 1. In particular, if threads access 16-byte words, at least two memory transactions are issued.

Unused words in a memory transaction are still read, so they waste bandwidth. To reduce waste, hardware will automatically issue the smallest memory transaction that contains the requested words. For example, if all the requested words lie in one half of a 128-byte segment, a 64-byte transaction will be issued.

More precisely, the following protocol is used to issue a memory transaction for a half-warp:

- ❑ Find the memory segment that contains the address requested by the lowest numbered active thread. Segment size is 32 bytes for 1-byte data, 64 bytes for 2-byte data, 128 bytes for 4-, 8- and 16-bit data.
- ❑ Find all other active threads whose requested address lies in the same segment.
- ❑ Reduce the transaction size, if possible:
 - ❑ If the transaction size is 128 bytes and only the lower or upper half is used, reduce the transaction size to 64 bytes;
 - ❑ If the transaction size is 64 bytes and only the lower or upper half is used, reduce the transaction size to 32 bytes.
- ❑ Carry out the transaction and mark the serviced threads as inactive.
- ❑ Repeat until all threads in the half-warp are serviced.

Figure 5-4 shows some examples of global memory accesses for devices of compute capability 1.2 and higher.



Left: coalesced `float` memory access, resulting in a single memory transaction.

Right: coalesced `float` memory access (divergent warp), resulting in a single memory transaction.

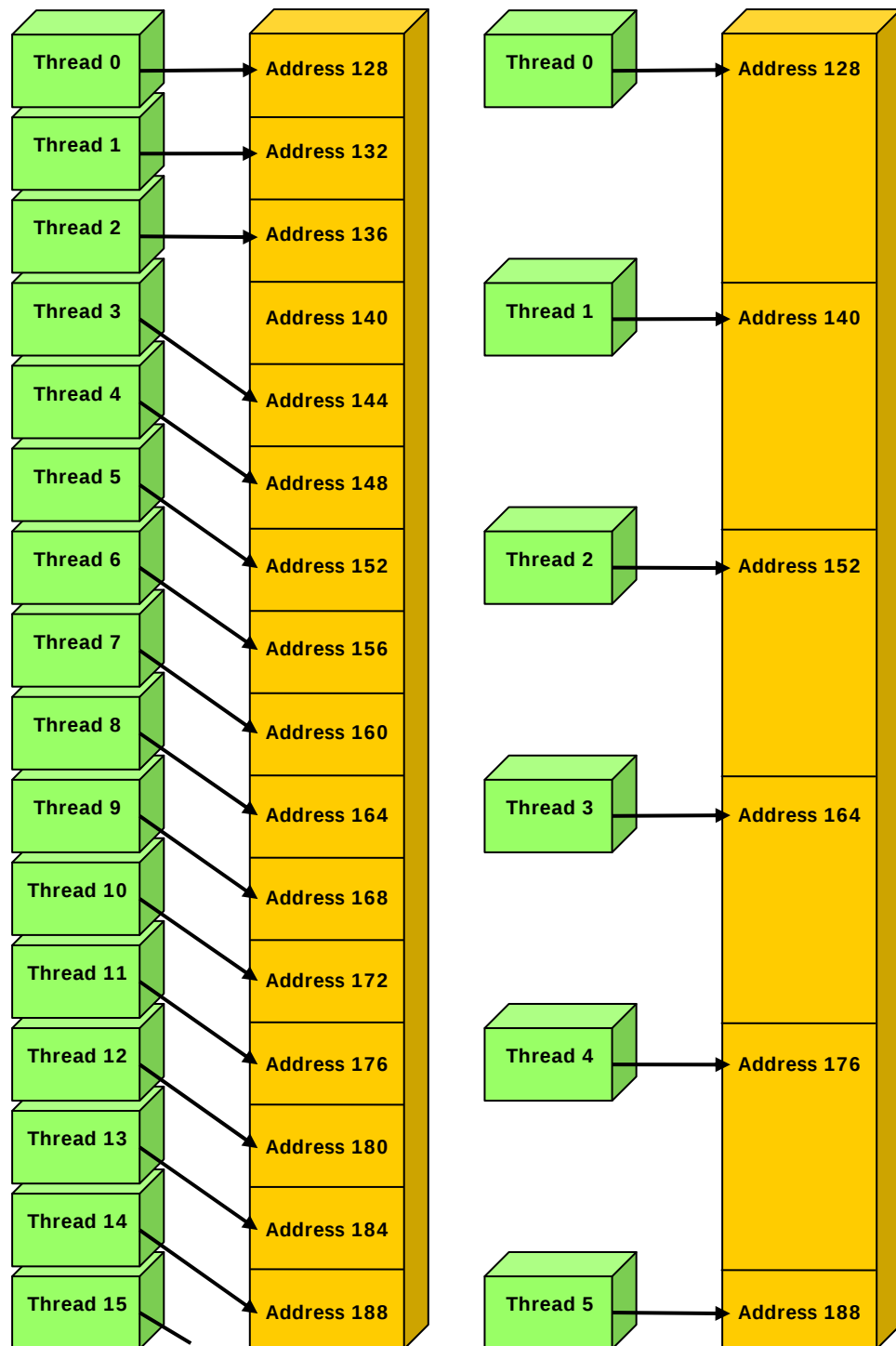
Figure 5-1. Examples of Coalesced Global Memory Access Patterns



Left: non-sequential `float` memory access, resulting in 16 memory transactions.

Right: access with a misaligned starting address, resulting in 16 memory transactions.

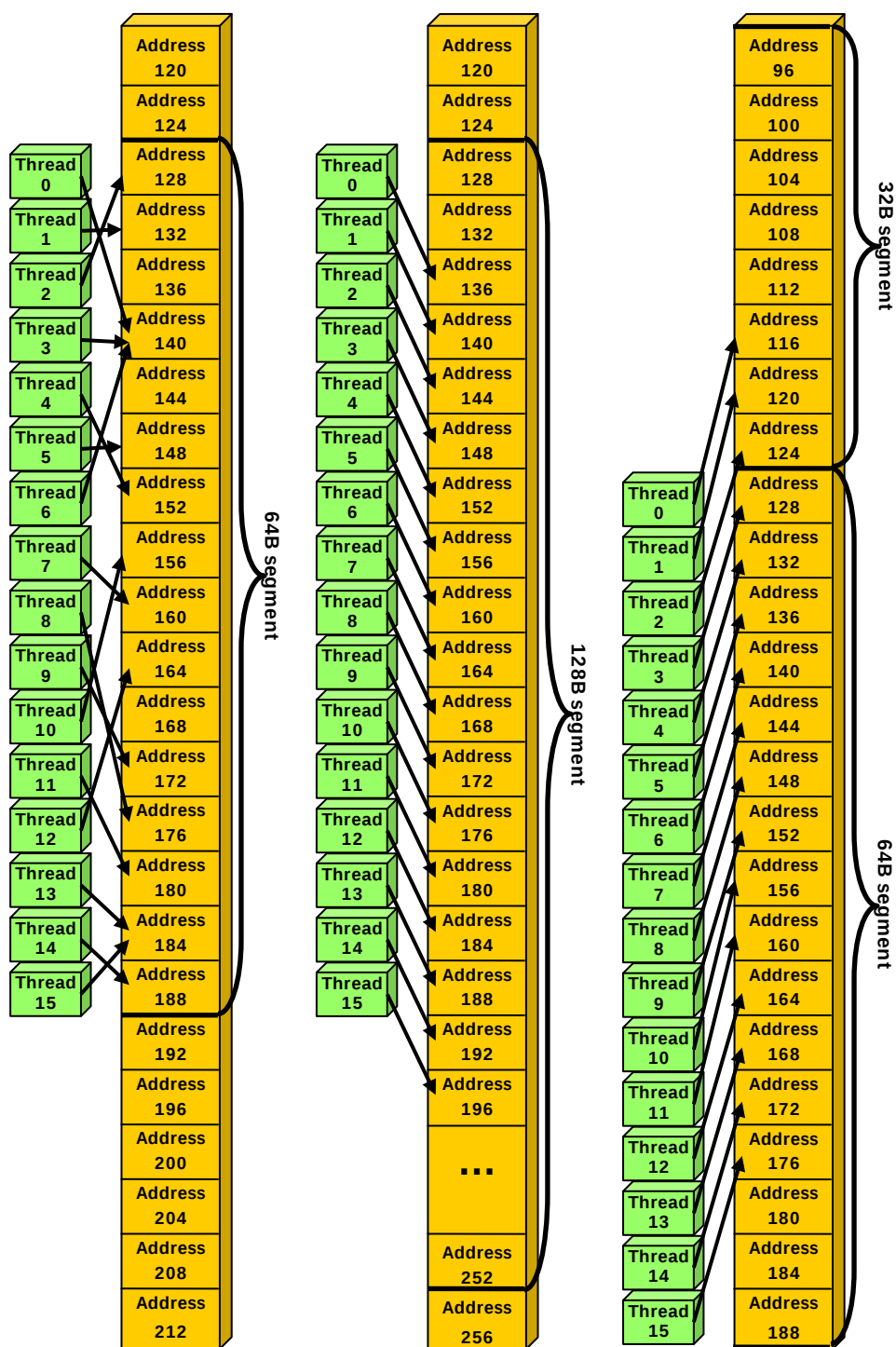
Figure 5-2. Examples of Global Memory Access Patterns That Are Non-Coalesced for Devices of Compute Capability 1.0 or 1.1



Left: non-contiguous `float` memory access, resulting in 16 memory transactions.

Right: non-coalesced `float3` memory access, resulting in 16 memory transactions.

Figure 5-3. Examples of Global Memory Access Patterns That Are Non-Coalesced for Devices of Compute Capability 1.0 or 1.1



Left: random **float** memory access within a 64B segment, resulting in one memory transaction.

Center: misaligned **float** memory access, resulting in one transaction.

Right: misaligned **float** memory access, resulting in two transactions.

Figure 5-4. Examples of Global Memory Access by Devices with Compute Capability 1.2 and Higher

Common Access Patterns

Array of Structures

A common global memory access pattern is when each thread of thread ID **tid** accesses one element of an array located at address **BaseAddress** of type **type*** using the following address:

$$\text{BaseAddress} + \text{tid}$$

To get memory coalescing, **type** must meet the size and alignment requirements discussed above. In particular, this means that if **type** is a structure larger than 16 bytes, it should be split into several structures that meet these requirements and the data should be laid out in memory as a list of several arrays of these structures instead of a single array of type **type***.

Two-Dimensional Array

Another common global memory access pattern is when each thread of index **(tx, ty)** accesses one element of a 2D array located at address **BaseAddress** of type **type*** and of width **width** using the following address:

$$\text{BaseAddress} + \text{width} * \text{ty} + \text{tx}$$

In such a case, one gets memory coalescing for all half-warps of the thread block only if:

- ❑ The width of the thread block is a multiple of half the warp size;
- ❑ **width** is a multiple of 16.

In particular, this means that an array whose width is not a multiple of 16 will be accessed much more efficiently if it is actually allocated with a width rounded up to the closest multiple of 16 and its rows padded accordingly. The **cudaMallocPitch()** and **cuMemAllocPitch()** functions and associated memory copy functions described in the reference manual enable programmers to write non-hardware-dependent code to allocate arrays that conform to these constraints.

5.1.2.2 Local Memory

Like the global memory space, the local memory space is not cached, so accesses to local memory are as expensive as accesses to global memory. Local memory accesses are always coalesced though since they are per-thread by definition.

Local memory accesses only occur for some automatic variables as mentioned in Section B.2.5. Inspection of the *PTX* assembly code (obtained by compiling with the **-ptx** or **-keep** option) will tell if a variable has been placed in local memory during the first compilation phases as it will be declared using the **.local** mnemonic and accessed using the **ld.local** and **st.local** mnemonics. If it has not, subsequent compilation phases might still decide otherwise though if they find it consumes too much register space for the targeted architecture. There is no way to check this for a particular variable, but the compiler reports total local memory usage per kernel (**lmem**) when compiling with the **--ptxas-options=-v** option.

Automatic variables that are likely to be placed in local memory are large structures or arrays that would consume too much register space, and arrays for which the compiler cannot determine that they are indexed with constant quantities.

5.1.2.3 Constant Memory

The constant memory space is cached so a read from constant memory costs one memory read from device memory only on a cache miss, otherwise it just costs one read from the constant cache.

For all threads of a half-warp, reading from the constant cache is as fast as reading from a register as long as all threads read the same address. The cost scales linearly with the number of different addresses read by all threads. We recommend having all threads of the entire warp read the same address as opposed to all threads within each of its halves only, as future devices will require it for full speed read.

5.1.2.4 Texture Memory

The texture memory space is cached so a texture fetch costs one memory read from device memory only on a cache miss, otherwise it just costs one read from the texture cache. The texture cache is optimized for 2D spatial locality, so threads of the same warp that read texture addresses that are close together will achieve best performance. Also, it is designed for streaming fetches with a constant latency, i.e. a cache hit reduces DRAM bandwidth demand, but not fetch latency.

Reading device memory through texture fetching present some benefits that can make it an advantageous alternative to reading device memory from global or constant memory:

- ❑ If the memory reads do not follow the access patterns that global or constant memory reads must respect to get good performance (see Sections 5.1.2.1 and 5.1.2.3), higher bandwidth can be achieved providing that there is locality in the texture fetches;
- ❑ The latency of addressing calculations is hidden better, possibly improving performance for applications that perform random accesses to the data;
- ❑ Packed data may be broadcast to separate variables in a single operation;
- ❑ 8-bit and 16-bit integer input data may be optionally converted to 32-bit floating-point values in the range [0.0, 1.0] or [-1.0, 1.0] (see Section 3.2.4.1).

However, within the same kernel call, the texture cache is not kept coherent with respect to global memory writes, so that any texture fetch to an address that has been written to via a global write in the same kernel call returns undefined data. In other words, a thread can safely read via texture some memory location only if this memory location has been updated by a previous kernel call or memory copy, but not if it has been previously updated by the same thread or another thread from the same kernel call. This is only relevant when fetching from linear memory as a kernel cannot write to CUDA arrays anyway.

5.1.2.5 Shared Memory

Because it is on-chip, the shared memory space is much faster than the local and global memory spaces. In fact, for all threads of a warp, accessing the shared memory is as fast as accessing a register as long as there are no bank conflicts between the threads, as detailed below.

To achieve high memory bandwidth, shared memory is divided into equally-sized memory modules, called banks, which can be accessed simultaneously. So, any memory read or write request made of n addresses that fall in n distinct memory banks can be serviced simultaneously, yielding an effective bandwidth that is n times as high as the bandwidth of a single module.

However, if two addresses of a memory request fall in the same memory bank, there is a bank conflict and the access has to be serialized. The hardware splits a memory request with bank conflicts into as many separate conflict-free requests as necessary, decreasing the effective bandwidth by a factor equal to the number of separate memory requests. If the number of separate memory requests is n , the initial memory request is said to cause n -way bank conflicts.

To get maximum performance, it is therefore important to understand how memory addresses map to memory banks in order to schedule the memory requests so as to minimize bank conflicts.

In the case of the shared memory space, the banks are organized such that successive 32-bit words are assigned to successive banks and each bank has a bandwidth of 32 bits per two clock cycles.

For devices of compute capability 1.x, the warp size is 32 and the number of banks is 16 (see Section 5.1); a shared memory request for a warp is split into one request for the first half of the warp and one request for the second half of the warp. As a consequence, there can be no bank conflict between a thread belonging to the first half of a warp and a thread belonging to the second half of the same warp.

A common case is for each thread to access a 32-bit word from an array indexed by the thread ID **tid** and with some stride **s**:

```
__shared__ float shared[32];
float data = shared[BaseIndex + s * tid];
```

In this case, the threads **tid** and **tid+n** access the same bank whenever **s*n** is a multiple of the number of banks **m** or equivalently, whenever **n** is a multiple of **m/d** where **d** is the greatest common divisor of **m** and **s**. As a consequence, there will be no bank conflict only if half the warp size is less than or equal to **m/d**. For devices of compute capability 1.x, this translates to no bank conflict only if **d** is equal to 1, or in other words, only if **s** is odd since **m** is a power of two.

Figure 5-5 and Figure 5-6 show some examples of conflict-free memory accesses while Figure 5-7 shows some examples of memory accesses that cause bank conflicts.

Other cases worth mentioning are when each thread accesses an element that is smaller or larger than 32 bits in size. For example, there are bank conflicts if an array of **char** is accessed the following way:

```
__shared__ char shared[32];
char data = shared[BaseIndex + tid];
```

because **shared[0]**, **shared[1]**, **shared[2]**, and **shared[3]**, for example, belong to the same bank. There are no bank conflicts however, if the same array is accessed the following way:

```
char data = shared[BaseIndex + 4 * tid];
```

There are also 2-way bank conflicts for arrays of **double**:

```
__shared__ double shared[32];
double data = shared[BaseIndex + tid];
```

since the memory request is compiled into two separate 32-bit requests. One way to avoid bank conflicts in this case is to split the **double** operands like in the following sample code:

```
__shared__ int shared_lo[32];
```

```

__shared__ int shared_hi[32];

double dataIn;
shared_lo[BaseIndex + tid] = __double2loint(dataIn);
shared_hi[BaseIndex + tid] = __double2hiint(dataIn);

double dataOut =
    __hiloInt2double(shared_hi[BaseIndex + tid],
                     shared_lo[BaseIndex + tid]);

```

It might not always improve performance though and will perform worse on future architectures.

A structure assignment is compiled into as many memory requests as necessary for each member in the structure, so the following code, for example:

```

__shared__ struct type shared[32];
struct type data = shared[BaseIndex + tid];

```

results in:

- Three separate memory reads without bank conflicts if **type** is defined as

```

struct type {
    float x, y, z;
};

```

since each member is accessed with a stride of three 32-bit words;

- Two separate memory reads with bank conflicts if **type** is defined as

```

struct type {
    float x, y;
};

```

since each member is accessed with a stride of two 32-bit words;

- Two separate memory reads with bank conflicts if **type** is defined as

```

struct type {
    float f;
    char c;
};

```

since each member is accessed with a stride of five bytes.

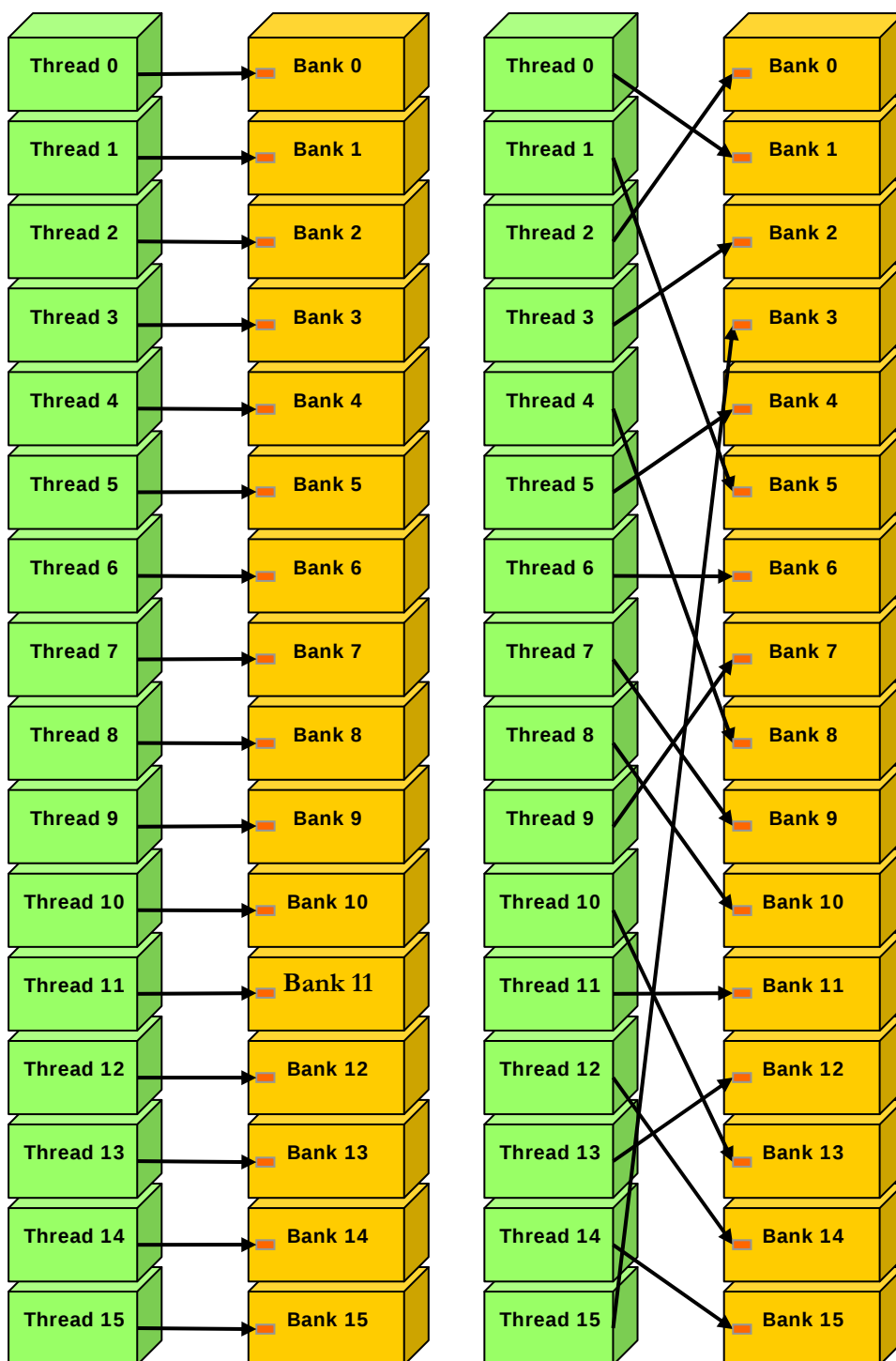
Finally, shared memory also features a broadcast mechanism whereby a 32-bit word can be read and broadcast to several threads simultaneously when servicing one memory read request. This reduces the number of bank conflicts when several threads of a half-warp read from an address within the same 32-bit word. More precisely, a memory read request made of several addresses is serviced in several steps over time – one step every two clock cycles – by servicing one conflict-free subset of these addresses per step until all addresses have been serviced; at each step, the subset is built from the remaining addresses that have yet to be serviced using the following procedure:

- Select one of the words pointed to by the remaining addresses as the broadcast word,
- Include in the subset:
 - All addresses that are within the broadcast word,
 - One address for each bank pointed to by the remaining addresses.

Which word is selected as the broadcast word and which address is picked up for each bank at each cycle are unspecified.

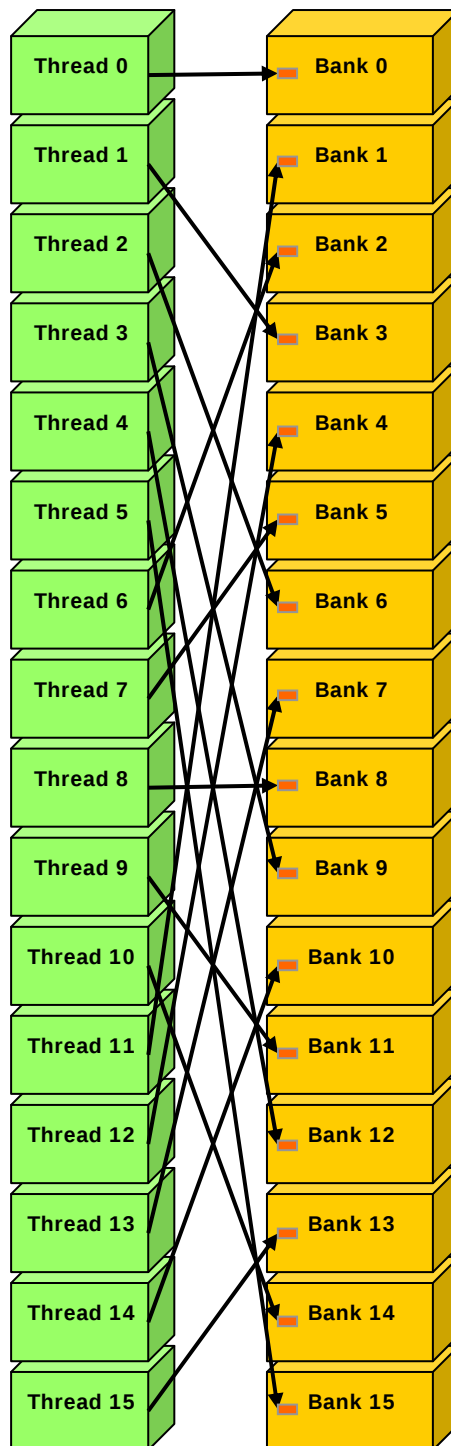
A common conflict-free case is when all threads of a half-warp read from an address within the same 32-bit word.

Figure 5-8 shows some examples of memory read accesses that involve the broadcast mechanism.



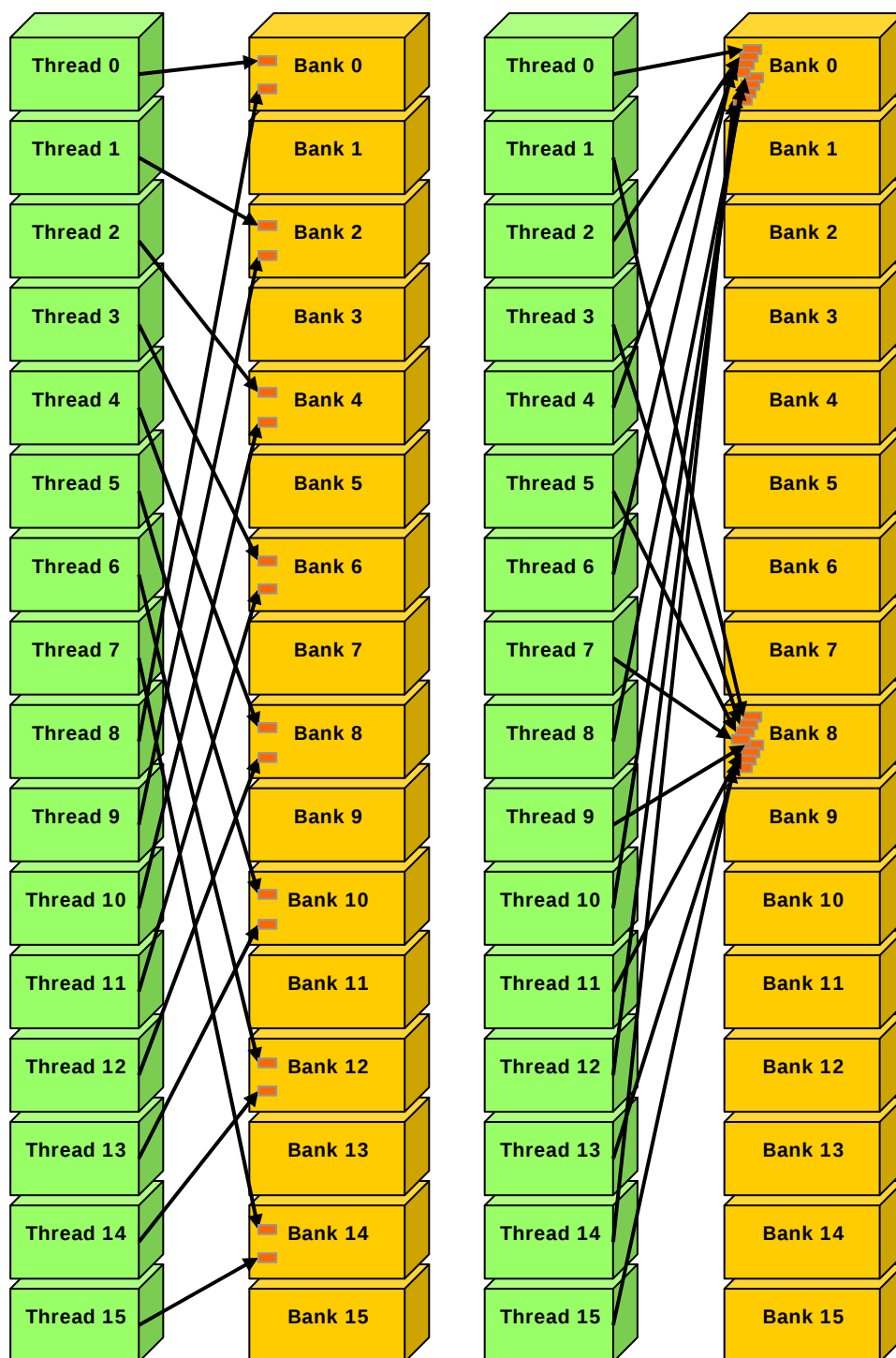
Left: linear addressing with a stride of one 32-bit word.
 Right: random permutation.

Figure 5-5. Examples of Shared Memory Access Patterns without Bank Conflicts



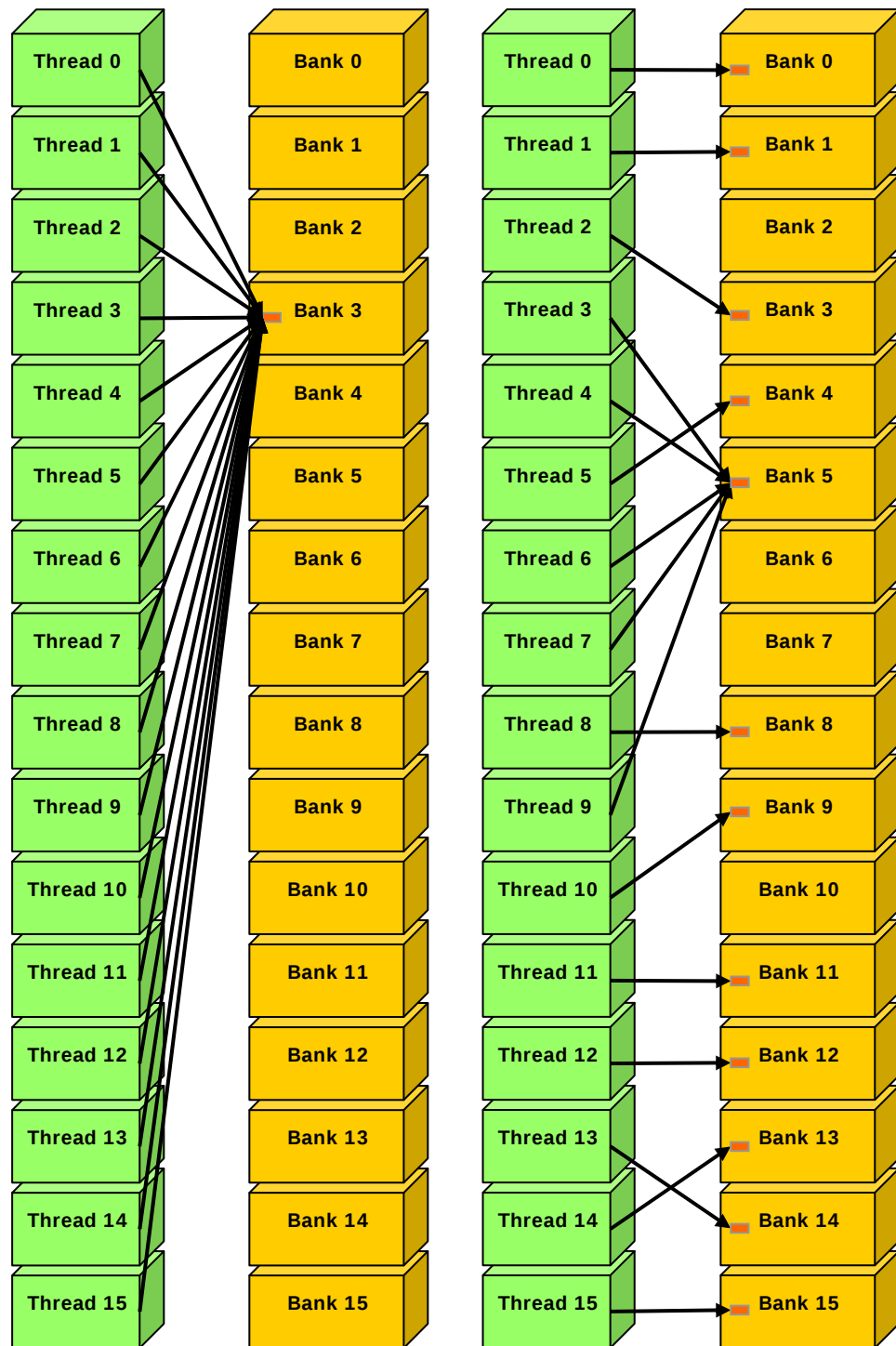
Linear addressing with a stride of three 32-bit words.

Figure 5-6. Example of a Shared Memory Access Pattern without Bank Conflicts



Left: Linear addressing with a stride of two 32-bit words causes 2-way bank conflicts.
 Right: Linear addressing with a stride of eight 32-bit words causes 8-way bank conflicts.

Figure 5-7. Examples of Shared Memory Access Patterns with Bank Conflicts



Left: This access pattern is conflict-free since all threads read from an address within the same 32-bit word.

Right: This access pattern causes either no bank conflicts if the word from bank 5 is chosen as the broadcast word during the first step or 2-way bank conflicts, otherwise.

Figure 5-8. Example of Shared Memory Read Access Patterns with Broadcast

5.1.2.6 Registers

Generally, accessing a register is zero extra clock cycles per instruction, but delays may occur due to register read-after-write dependencies and register memory bank conflicts.

The delays introduced by read-after-write dependencies can be ignored as soon as there are at least 192 active threads per multiprocessor to hide them.

The compiler and thread scheduler schedule the instructions as optimally as possible to avoid register memory bank conflicts. They achieve best results when the number of threads per block is a multiple of 64. Other than following this rule, an application has no direct control over these bank conflicts. In particular, there is no need to pack data into **float4** or **int4** types.

5.2 Execution Configuration

How the execution configuration affects the execution time of a kernel launch generally depends on the kernel code. Experimentation is therefore recommended. There are however general guidelines, described in this section.

For a start, the kernel will simply fail to launch if the number of threads per block either is above the maximum number of threads per block as specified in Appendix A, or requires too many registers or shared memory than available per multiprocessor as mentioned in Section 4.1. The total number of registers required for a block is equal to

$$\text{ceil}(R \times \text{ceil}(T, 32), \frac{R_{\max}}{32})$$

where R is the number of registers required for the kernel, $R_{\max}$ is the number of registers per multiprocessor given in Appendix A, T is the number of threads per block, and $\text{ceil}(x, y)$ is equal to x rounded up to the nearest multiple of y . The total amount of shared memory required for a block is equal to the sum of the amount of statically allocated shared memory, the amount of dynamically allocated shared memory, and the amount of shared memory used to pass the kernel's arguments. The number of registers a kernel compiles to and the local, shared, and constant memory usages are reported by the compiler when compiling with the **--ptxas-options=-v** option. Note that each **double** or **long long** variable uses two registers. However, devices of compute capability 1.2 and higher have twice as many registers per multiprocessor as devices with lower compute capability.

Then, given a total number of threads per grid, the number of threads per block might be dictated by the need to have enough blocks in the grid to maximize the utilization of the available computing resources. First, there should be at least as many blocks as there are multiprocessors in the device. Then, running only one block per multiprocessor will force the multiprocessor to idle during thread synchronization and also during device memory reads if there are not enough threads per block to cover the load latency. It is therefore usually better to allow for two or more blocks to be active on each multiprocessor to allow overlap between blocks that wait and blocks that can run. For this to happen, not only should there be at least twice as many blocks as there are multiprocessors in the device, but also the amount of registers and shared memory required per block must be low enough

to allow for more than one active block (see Section 4.1). More thread blocks stream in pipeline fashion through the device and amortize overhead even more. The number of blocks per grid should be at least 100 if one wants it to scale to future devices; 1000 blocks will scale across several generations.

With a high enough number of blocks, the number of threads per block should be chosen as a multiple of the warp size to avoid wasting computing resources with under-populated warps, or better, a multiple of 64 for the reason invoked in Section 5.1.2.6. Allocating more threads per block is better for efficient time slicing, but the more threads per block, the fewer registers are available per thread, which might prevent the kernel invocation from succeeding.

Usually, 64 threads per block is minimal and makes sense only if there are multiple active blocks per multiprocessor; 192 or 256 threads per block is better and usually allows for enough registers to compile.

The ratio of the number of active warps per multiprocessor to the maximum number of active warps (given in Appendix A) is called the multiprocessor *occupancy*. In order to maximize occupancy, the compiler attempts to minimize register usage while keeping the number of instructions and local memory usage to a minimum. This can be controlled using the **-maxrregcount** compiler option. The CUDA Software Development Kit provides a spreadsheet to assist programmers in choosing thread block size based on shared memory and register requirements.

5.3 Data Transfer between Host and Device

The bandwidth between device memory and the device is much higher than the bandwidth between device memory and host memory. Therefore, one should strive to minimize data transfer between the host and the device, for example, by moving more code from the host to the device, even if that means running kernels with low parallelism computations. Intermediate data structures may be created in device memory, operated on by the device, and destroyed without ever being mapped by the host or copied to host memory.

Also, because of the overhead associated with each transfer, batching many small transfers into a big one always performs much better than making each transfer separately.

Finally, higher performance for data transfers between host and device is achieved by using page-locked host memory as described in Section 3.2.5.

In addition, when using mapped page-locked memory (Section 3.2.5.3), there is no need to allocate any device memory and to explicitly copy data between device and host memory. Data transfers are implicitly performed each time the kernel accesses the mapped memory. For maximum performance, these memory accesses must be coalesced like if they were accesses to global memory (see Section 5.1.2.1).

Assuming that they are and that the mapped memory is read or written only once, using mapped page-locked memory instead of explicit copies between device and host memory can be a win performance-wise.

On integrated systems where device memory and host memory are physically the same, any copy between host and device memory is superfluous and mapped page-locked memory should be used instead. Applications may query whether a device is

integrated or not by calling `cudaGetDeviceProperties()` and checking the **integrated** property or checking the `CU_DEVICE_ATTRIBUTE_INTEGRATED` attribute using `cuDeviceGetAttribute()`.

5.4 Warp-Level Synchronization

Because a warp executes one common instruction at a time, threads within a warp are implicitly synchronized and this can be used to omit `__syncthreads()` for better performance.

In the following code sample, for example, both calls to `__syncthreads()` are required to get the expected result (i.e. `result[i] = 2 * myArray[i]` for `i > 0`). Without synchronization, any of the two references to `myArray[tid]` could return either 2 or the value initially stored in `myArray`, depending on whether the memory read occurs before or after the memory write from `myArray[tid + 1] = 2`.

```
// myArray is an array of integers located in global or shared
// memory
__global__ void myKernel(int* result) {
    int tid = threadIdx.x;
    ...
    int ref1 = myArray[tid] * 1;
    __syncthreads();
    myArray[tid + 1] = 2;
    __syncthreads();
    int ref2 = myArray[tid] * 1;
    result[tid] = ref1 * ref2;
    ...
}
```

However, in the following slightly modified code sample, threads are guaranteed to belong to the same warp, so that there is no need for any `__syncthreads()`.

```
// myArray is an array of integers located in global or shared
// memory
__global__ void myKernel(int* result) {
    int tid = threadIdx.x;
    ...
    if (tid < warpSize) {
        int ref1 = myArray[tid] * 1;
        myArray[tid + 1] = 2;
        int ref2 = myArray[tid] * 1;
        result[tid] = ref1 * ref2;
    }
    ...
}
```

Simply removing the `__syncthreads()` is not enough however; `myArray` also needs to be declared as volatile as described in Section B.2.4.

5.5 Overall Performance Optimization Strategies

Performance optimization revolves around three basic strategies:

- ❑ Maximizing parallel execution;
- ❑ Optimizing memory usage to achieve maximum memory bandwidth;
- ❑ Optimizing instruction usage to achieve maximum instruction throughput.

Maximizing parallel execution starts with structuring the algorithm in a way that exposes as much data parallelism as possible. At points in the algorithm where parallelism is broken because some threads need to synchronize in order to share data between each other, there are two cases: Either these threads belong to the same block, in which case they should use `__syncthreads()` and share data through shared memory within the same kernel call, or they belong to different blocks, in which case they must share data through global memory using two separate kernel invocations, one for writing to and one for reading from global memory.

Once the parallelism of the algorithm has been exposed it needs to be mapped to the hardware as efficiently as possible. This is done by carefully choosing the execution configuration of each kernel invocation as detailed in Section 5.2.

The application should also maximize parallel execution at a higher level by explicitly exposing concurrent execution on the device through streams, as described in Section 3.2.6.1, as well as maximizing concurrent execution between host and device.

Optimizing memory usage starts with minimizing data transfers with low-bandwidth. That means minimizing data transfers between the host and the device, as detailed in Section 5.3, since these have much lower bandwidth than data transfers between the device and global memory. That also means minimizing data transfers between the device and global memory by maximizing use of shared memory on the device, as mentioned in Section 5.1.2. Sometimes, the best optimization might even be to avoid any data transfer in the first place by simply recomputing the data instead whenever it is needed.

As detailed in Sections 5.1.2.1, 5.1.2.3, 5.1.2.4, and 5.1.2.5, the effective bandwidth can vary by an order of magnitude depending on access pattern for each type of memory. The next step in optimizing memory usage is therefore to organize memory accesses as optimally as possible based on the optimal memory access patterns. This optimization is especially important for global memory accesses as global memory bandwidth is low and its latency is hundreds of clock cycles (see Section 5.1.1.3). Shared memory accesses, on the other hand, are usually worth optimizing only in case they have a high degree of bank conflicts.

As for optimizing instruction usage, the use of arithmetic instructions with low throughput (see Section 5.1.1.1) should be minimized. This includes trading precision for speed when it does not affect the end result, such as using intrinsic instead of regular functions (intrinsic functions are listed in Section C.2) or single-precision instead of double-precision. Particular attention must be paid to control flow instructions due to the SIMT nature of the device as detailed in Section 5.1.1.2.

Appendix A.

Technical Specifications

A.1 General Specifications

The general specifications and features of a compute device depend on its compute capability (see Section 2.5).

The following sections describe the technical specifications and features associated to each compute capability. The specifications for a given compute capability are the same as for the compute capability just below unless otherwise mentioned. Similarly, any feature supported for a given compute capability is supported for any higher compute capability.

The compute capability and number of multiprocessors of all CUDA-enabled devices are given in the following table:

	Number of Multiprocessors (1 Multiprocessor = 8 Processors)	Compute Capability
GeForce GTX 295	2x30	1.3
GeForce GTX 285, GTX 280	30	1.3
GeForce GTX 260	24	1.3
GeForce 9800 GX2	2x16	1.1
GeForce GTS 250, GTS 150, 9800 GTX, 9800 GTX+, 8800 GTS 512	16	1.1
GeForce 8800 Ultra, 8800 GTX	16	1.0
GeForce 9800 GT, 8800 GT, GTX 280M, 9800M GTX	14	1.1
GeForce GT 130, 9600 GSO, 8800 GS, 8800M GTX, GTX 260M, 9800M GT	12	1.1
GeForce 8800 GTS	12	1.0
GeForce 9600 GT, 8800M GTS, 9800M GTS	8	1.1
GeForce 9700M GT	6	1.1
GeForce GT 120, 9500 GT, 8600 GTS, 8600 GT, 9700M GT, 9650M GS, 9600M GT, 9600M GS, 9500M GS, 8700M GT, 8600M GT, 8600M GS	4	1.1
GeForce G100, 8500 GT, 8400 GS, 8400M GT,	2	1.1

9500M G, 9300M G, 8400M GS, 9400 mGPU, 9300 mGPU, 8300 mGPU, 8200 mGPU, 8100 mGPU		
GeForce 9300M GS, 9200M GS, 9100M G, 8400M G	1	1.1
Tesla S1070	4x30	1.3
Tesla C1060	30	1.3
Tesla S870	4x16	1.0
Tesla D870	2x16	1.0
Tesla C870	16	1.0
Quadro Plex 2200 D2	2x30	1.3
Quadro Plex 2100 D4	4x14	1.1
Quadro Plex 2100 Model S4	4x16	1.0
Quadro Plex 1000 Model IV	2x16	1.0
Quadro FX 5800	30	1.3
Quadro FX 4800	24	1.3
Quadro FX 4700 X2	2x14	1.1
Quadro FX 3700M	16	1.1
Quadro FX 5600	16	1.0
Quadro FX 3700	14	1.1
Quadro FX 3600M	12	1.1
Quadro FX 4600	12	1.0
Quadro FX 2700M	6	1.1
Quadro FX 1700, FX 570, NVS 320M, FX 1700M, FX 1600M, FX 770M, FX 570M	4	1.1
Quadro FX 370, NVS 290, NVS 140M, NVS 135M, FX 360M	2	1.1
Quadro FX 370M, NVS 130M	1	1.1

The number of multiprocessors, the clock frequency and the total amount of device memory can be queried using the runtime (see reference manual).

A.1.1 Specifications for Compute Capability 1.0

- ❑ The maximum number of threads per block is 512;
- ❑ The maximum sizes of the x-, y-, and z-dimension of a thread block are 512, 512, and 64, respectively;
- ❑ The maximum size of each dimension of a grid of thread blocks is 65535;
- ❑ The warp size is 32 threads;
- ❑ The number of registers per multiprocessor is 8192;
- ❑ The amount of shared memory available per multiprocessor is 16 KB organized into 16 banks (see Section 5.1.2.5);
- ❑ The total amount of constant memory is 64 KB;
- ❑ The total amount of local memory per thread is 16 KB;
- ❑ The cache working set for constant memory is 8 KB per multiprocessor;

- ❑ The cache working set for texture memory varies between 6 and 8 KB per multiprocessor;
- ❑ The maximum number of active blocks per multiprocessor is 8;
- ❑ The maximum number of active warps per multiprocessor is 24;
- ❑ The maximum number of active threads per multiprocessor is 768;
- ❑ For a one-dimensional texture reference bound to a CUDA array, the maximum width is 2^{13} ;
- ❑ For a one-dimensional texture reference bound to linear memory, the maximum width is 2^{27} ;
- ❑ For a two-dimensional texture reference bound to linear memory or a CUDA array, the maximum width is 2^{16} and the maximum height is 2^{15} ;
- ❑ For a three-dimensional texture reference bound to a CUDA array, the maximum width is 2^{11} , the maximum height is 2^{11} , and the maximum depth is 2^{11} ;
- ❑ The limit on kernel size is 2 millions of microcode instructions;

A.1.2 Specifications for Compute Capability 1.1

- ❑ Support for atomic functions operating on 32-bit words in global memory (see Section B.10).

A.1.3 Specifications for Compute Capability 1.2

- ❑ Support for atomic functions operating in shared memory and atomic functions operating on 64-bit words in global memory (see Section B.10);
- ❑ Support for warp vote functions (see Section B.11);
- ❑ The number of registers per multiprocessor is 16384;
- ❑ The maximum number of active warps per multiprocessor is 32;
- ❑ The maximum number of active threads per multiprocessor is 1024.

A.1.4 Specifications for Compute Capability 1.3

- ❑ Support for double-precision floating-point numbers.

A.2 Floating-Point Standard

All compute devices follow the IEEE-754 standard for binary floating-point arithmetic with the following deviations:

- ❑ There is no dynamically configurable rounding mode; however, most of the operations support IEEE rounding modes, exposed via device functions;
- ❑ There is no mechanism for detecting that a floating-point exception has occurred and all operations behave as if the IEEE-754 exceptions are always masked, and deliver the masked response as defined by IEEE-754 if there is an exceptional event; for the same reason, while SNaN encodings are supported, they are not signaling;

- ❑ Absolute value and negation are not compliant with IEEE-754 with respect to NaNs; these are passed through unchanged;
- ❑ For single-precision floating-point numbers only:
 - ❑ Denormalized numbers are not supported; floating-point arithmetic and comparison instructions convert denormalized operands to zero prior to the floating-point operation;
 - ❑ Underflowed results are flushed to zero;
 - ❑ The result of an operation involving one or more input NaNs is the quiet NaN of bit pattern 0x7fffffff; note that;
 - ❑ Some instructions are not IEEE-compliant:
 - ❑ Addition and multiplication are often combined into a single multiply-add instruction (FMAD), which truncates the intermediate result of the multiplication;
 - ❑ Division is implemented via the reciprocal in a non-standard-compliant way;
 - ❑ Square root is implemented via the reciprocal square root in a non-standard-compliant way;
 - ❑ For addition and multiplication, only round-to-nearest-even and round-towards-zero are supported via static rounding modes; directed rounding towards +/- infinity is not supported;

But, IEEE-compliant software (and therefore slower) implementations are provided through the following intrinsics from Section C.2.1:

- ❑ **__fmaf_r{n,z,u,d}(float, float, float)**: single-precision fused multiply-add with IEEE rounding modes,
- ❑ **__frcp_r{n,z,u,d}(float)**: single-precision reciprocal with IEEE rounding modes,
- ❑ **__fdiv_r{n,z,u,d}(float, float)**: single-precision division with IEEE rounding modes,
- ❑ **__fsqrt_r{n,z,u,d}(float)**: single-precision square root with IEEE rounding modes,
- ❑ **__fadd_r[u,d](float, float)**: single-precision addition with IEEE directed rounding,
- ❑ **__fmul_r[u,d](float, float)**: single-precision multiplication with IEEE directed rounding;
- ❑ For double-precision floating-point numbers only:
 - ❑ Round-to-nearest-even is the only supported IEEE rounding mode for reciprocal, division, and square root.

In accordance to the IEEE-754R standard, if one of the input parameters to **fminf()**, **fmin()**, **fmaxf()**, or **fmax()** is NaN, but not the other, the result is the non-NaN parameter.

The conversion of a floating-point value to an integer value in the case where the floating-point value falls outside the range of the integer format is left undefined by IEEE-754. For compute devices, the behavior is to clamp to the end of the supported range. This is unlike the x86 architecture behaves.

Appendix B. C Extensions

B.1 Function Type Qualifiers

Function type qualifiers specify whether a function executes on the host or on the device and whether it is callable from the host or from the device.

B.1.1 `__device__`

The `__device__` qualifier declares a function that is:

- ❑ Executed on the device
- ❑ Callable from the device only.

B.1.2 `__global__`

The `__global__` qualifier declares a function as being a kernel. Such a function is:

- ❑ Executed on the device,
- ❑ Callable from the host only.

B.1.3 `__host__`

The `__host__` qualifier declares a function that is:

- ❑ Executed on the host,
- ❑ Callable from the host only.

It is equivalent to declare a function with only the `__host__` qualifier or to declare it without any of the `__host__`, `__device__`, or `__global__` qualifier; in either case the function is compiled for the host only.

However, the `__host__` qualifier can also be used in combination with the `__device__` qualifier, in which case the function is compiled for both the host and the device.

B.1.4 Restrictions

`__device__` and `__global__` functions do not support recursion.

`__device__` and `__global__` functions cannot declare static variables inside their body.

`__device__` and `__global__` functions cannot have a variable number of arguments.

`__device__` functions cannot have their address taken; function pointers to `__global__` functions, on the other hand, are supported.

The `__global__` and `__host__` qualifiers cannot be used together.

`__global__` functions must have void return type.

Any call to a `__global__` function must specify its execution configuration as described in Section B.12.

A call to a `__global__` function is asynchronous, meaning it returns before the device has completed its execution.

`__global__` function parameters are currently passed via shared memory to the device and limited to 256 bytes.

B.2 Variable Type Qualifiers

Variable type qualifiers specify the memory location on the device of a variable.

B.2.1 `__device__`

The `__device__` qualifier declares a variable that resides on the device.

At most one of the other type qualifiers defined in the next three sections may be used together with `__device__` to further specify which memory space the variable belongs to. If none of them is present, the variable:

- ❑ Resides in global memory space,
- ❑ Has the lifetime of an application,
- ❑ Is accessible from all the threads within the grid and from the host through the runtime library.

B.2.2 `__constant__`

The `__constant__` qualifier, optionally used together with `__device__`, declares a variable that:

- ❑ Resides in constant memory space,
- ❑ Has the lifetime of an application,
- ❑ Is accessible from all the threads within the grid and from the host through the runtime library.

B.2.3 `__shared__`

The `__shared__` qualifier, optionally used together with `__device__`, declares a variable that:

- ❑ Resides in the shared memory space of a thread block,
- ❑ Has the lifetime of the block,
- ❑ Is only accessible from all the threads within the block.

When declaring a variable in shared memory as an external array such as

```
extern __shared__ float shared[];
```

the size of the array is determined at launch time (see Section B.12). All variables declared in this fashion, start at the same address in memory, so that the layout of the variables in the array must be explicitly managed through offsets. For example, if one wants the equivalent of

```
short array0[128];
float array1[64];
int array2[256];
```

in dynamically allocated shared memory, one could declare and initialize the arrays the following way:

```
extern __shared__ char array[];
__device__ void func() // __device__ or __global__ function
{
    short* array0 = (short*)array;
    float* array1 = (float*)&array0[128];
    int* array2 = (int*)&array1[64];
}
```

Note that pointers need to be aligned to the type they point to, so the following code, for example, does not work since **array1** is not aligned to 4 bytes.

```
extern __shared__ char array[];
__device__ void func() // __device__ or __global__ function
{
    short* array0 = (short*)array;
    float* array1 = (float*)&array0[127];
}
```

Alignment requirements for the built-in vector types are listed in Table B-1.

B.2.4 Volatile

Only after the execution of a `__threadfence_block()`, `__threadfence()`, or `__syncthreads()` (Sections B.5 and B.6) are prior writes to global or shared memory guaranteed to be visible by other threads. As long as this requirement is met, the compiler is free to optimize reads and writes to global or shared memory. For example, in the code sample below, the first reference to **myArray[tid]** compiles into a global or shared memory read instruction, but the second reference does not as the compiler simply reuses the result of the first read.

```
// myArray is an array of non-zero integers
// located in global or shared memory
__global__ void myKernel(int* result) {
    int tid = threadIdx.x;
    int ref1 = myArray[tid] * 1;
```

```

    myArray[tid + 1] = 2;
    int ref2 = myArray[tid] * 1;
    result[tid] = ref1 * ref2;
}

```

Therefore, **ref2** cannot possibly be equal to **2** in thread **tid** as a result of thread **tid-1** overwriting **myArray[tid]** by 2.

This behavior can be changed using the **volatile** keyword: If a variable located in global or shared memory is declared as volatile, the compiler assumes that its value can be changed at any time by another thread and therefore any reference to this variable compiles to an actual memory read instruction.

Note that even if **myArray** is declared as volatile in the code sample above, there is no guarantee, in general, that **ref2** will be equal to 2 in thread **tid** since thread **tid** might read **myArray[tid]** into **ref2** before thread **tid-1** overwrites its value by 2. Synchronization is required as mentioned in Section 5.4.

B.2.5 Restrictions

These qualifiers are not allowed on **struct** and **union** members, on formal parameters and on local variables within a function that executes on the host.

__shared__ and **__constant__** variables have implied static storage.

__device__, **__shared__** and **__constant__** variables cannot be defined as external using the **extern** keyword. The only exception is for dynamically allocated **__shared__** variables as described in Section B.2.3.

__device__ and **__constant__** variables are only allowed at file scope.

__constant__ variables cannot be assigned to from the device, only from the host through host runtime functions (Sections 3.2.1 and 3.3.4).

__shared__ variables cannot have an initialization as part of their declaration.

An automatic variable declared in device code without any of these qualifiers generally resides in a register. However in some cases the compiler might choose to place it in local memory, which can have adverse performance consequences as detailed in Section 5.1.2.2.

Pointers in code that is executed on the device are supported as long as the compiler is able to resolve whether they point to either the shared memory space or the global memory space, otherwise they are restricted to only point to memory allocated or declared in the global memory space.

Dereferencing a pointer either to global or shared memory in code that is executed on the host or to host memory in code that is executed on the device results in an undefined behavior, most often in a segmentation fault and application termination.

The address obtained by taking the address of a **__device__**, **__shared__** or **__constant__** variable can only be used in device code. The address of a **__device__** or **__constant__** variable obtained through **cudaGetSymbolAddress()** as described in Section 3.3.4 can only be used in host code.

B.3 Built-in Vector Types

B.3.1 char1, uchar1, char2, uchar2, char3, uchar3, char4, uchar4, short1, ushort1, short2, ushort2, short3, ushort3, short4, ushort4, int1, uint1, int2, uint2, int3, uint3, int4, uint4, long1, ulong1, long2, ulong2, long3, ulong3, long4, ulong4, longlong1, longlong2, float1, float2, float3, float4, double1, double2

These are vector types derived from the basic integer and floating-point types. They are structures and the 1st, 2nd, 3rd, and 4th components are accessible through the fields **x**, **y**, **z**, and **w**, respectively. They all come with a constructor function of the form **make_<type name>**; for example,

```
int2 make_int2(int x, int y);
```

which creates a vector of type **int2** with value (**x**, **y**).

In host code, the alignment requirement of a vector type is equal to the alignment requirement of its base type. This is not always the case in device code as detailed in Table B-1.

Table B-1. Alignment Requirements in Device Code

Type	Alignment
char1, uchar1	1
char2, uchar2	2
char3, uchar3	1
char4, uchar4	4
short1, ushort1	2
short2, ushort2	4
short3, ushort3	2
short4, ushort4	8
int1, uint1	4
int2, uint2	8
int3, uint3	4
int4, uint4	16
long1, ulong1	Same as int1 or longlong1 depending on platform
long2, ulong2	Same as int2 or longlong2 depending on platform
long3, ulong3	Same as int3 depending on platform
long4, ulong4	Same as int4 depending on platform
longlong1	8
longlong2	16
float1	4

float2	8
float3	4
float4	16
double1	8
double2	16

B.3.2 dim3

This type is an integer vector type based on **uint3** that is used to specify dimensions. When defining a variable of type **dim3**, any component left unspecified is initialized to 1.

B.4 Built-in Variables

Built-in variables specify the grid and block dimensions and the block and thread indices. They are only valid within functions that are executed on the device.

B.4.1 gridDim

This variable is of type **dim3** (see Section B.3.2) and contains the dimensions of the grid.

B.4.2 blockDim

This variable is of type **uint3** (see Section B.3.1) and contains the block index within the grid.

B.4.3 blockDim

This variable is of type **dim3** (see Section B.3.2) and contains the dimensions of the block.

B.4.4 threadIdx

This variable is of type **uint3** (see Section B.3.1) and contains the thread index within the block.

B.4.5 warpSize

This variable is of type **int** and contains the warp size in threads (see Section 4.1 for the definition of a warp).

B.4.6 Restrictions

- ❑ It is not allowed to take the address of any of the built-in variables.

- It is not allowed to assign values to any of the built-in variables.

B.5 Memory Fence Functions

```
void __threadfence();
```

waits until all global and shared memory accesses made by the calling thread prior to **__threadfence()** are visible to all threads in the device for global memory accesses and all threads in the thread block for shared memory accesses.

```
void __threadfence_block();
```

waits until all global and shared memory accesses made by the calling thread prior to **__threadfence_block()** are visible to all threads in the thread block.

In general, when a thread issues a series of writes to memory in a particular order, other threads may see the effects of these memory writes in a different order.

__threadfence() and **__threadfence_block()** can be used to enforce some ordering.

One use case is when threads consume some data produced by other threads as illustrated by the following code sample of a kernel that computes the sum of an array of N numbers in one call. Each block first sums a subset of the array and stores the result in global memory. When all blocks are done, the last block done reads each of these partial sums from global memory and sums them to obtain the final result. In order to determine which block is finished last, each block atomically increments a counter to signal that it is done with computing and storing its partial sum (see Section B.10 about atomic functions). The last block is the one that receives the counter value equal to **gridDim.x-1**. If no fence is placed between storing the partial sum and incrementing the counter, the counter might increment before the partial sum is stored and therefore, might reach **gridDim.x-1** and let the last block start reading partial sums before they have been actually updated in memory.

```
__device__ unsigned int count = 0;
__shared__ bool isLastBlockDone;
__global__ void sum(const float* array, unsigned int N,
                  float* result)
{
    // Each block sums a subset of the input array
    float partialSum = calculatePartialSum(array, N);

    if (threadIdx.x == 0) {
        // Thread 0 of each block stores the partial sum
        // to global memory
        result[blockIdx.x] = partialSum;

        // Thread 0 makes sure its result is visible to
        // all other threads
        __threadfence();

        // Thread 0 of each block signals that it is done
        unsigned int value = atomicInc(&count, blockDim.x);

        // Thread 0 of each block determines if its block is
```

```

        // the last block to be done
        isLastBlockDone = (value == (gridDim.x - 1));
    }

    // Synchronize to make sure that each thread reads
    // the correct value of isLastBlockDone
    __syncthreads();

    if (isLastBlockDone) {

        // The last block sums the partial sums
        // stored in result[0 .. gridDim.x-1]
        float totalSum = calculateTotalSum(result);

        if (threadIdx.x == 0) {

            // Thread 0 of last block stores total sum
            // to global memory and resets count so that
            // next kernel call works properly
            result[0] = totalSum;
            count = 0;
        }
    }
}

```

B.6 Synchronization Function

```
void __syncthreads();
```

waits until all threads in the thread block have reached this point and all global and shared memory accesses made by these threads prior to **__syncthreads()** are visible to all threads in the block.

__syncthreads() is used to coordinate communication between the threads of the same block. When some threads within a block access the same addresses in shared or global memory, there are potential read-after-write, write-after-read, or write-after-write hazards for some of these memory accesses. These data hazards can be avoided by synchronizing threads in-between these accesses.

__syncthreads() is allowed in conditional code but only if the conditional evaluates identically across the entire thread block, otherwise the code execution is likely to hang or produce unintended side effects.

B.7 Mathematical Functions

Section C.1 contains a comprehensive list of the C/C++ standard library mathematical functions that are currently supported in device code, along with their respective error bounds. When executed in host code, a given function uses the C runtime implementation if available.

For some of the functions of Section C.1, a less accurate, but faster version exists in the device runtime component; it has the same name prefixed with **__** (such as

`__sinf(x)`). These intrinsic functions are listed in Section C.2, along with their respective error bounds.

The compiler has an option (`-use_fast_math`) that forces each function in Table B-2 to compile to its intrinsic counterpart. In addition to reduce accuracy of the affected functions, it may also cause some differences in special case handling. A more robust approach is to selectively replace mathematical function calls by calls to intrinsic functions only where it is merited by the performance gains and where changed properties such as reduced accuracy and different special case handling can be tolerated.

Table B-2. Functions Affected by `-use_fast_math`

Operator/Function	Device Function
<code>x/y</code>	<code>__fdividef(x,y)</code>
<code>sinf(x)</code>	<code>__sinf(x)</code>
<code>cosf(x)</code>	<code>__cosf(x)</code>
<code>tanf(x)</code>	<code>__tanf(x)</code>
<code>sincosf(x, sptr, cptr)</code>	<code>__sincosf(x, sptr, cptr)</code>
<code>logf(x)</code>	<code>__logf(x)</code>
<code>log2f(x)</code>	<code>__log2f(x)</code>
<code>log10f(x)</code>	<code>__log10f(x)</code>
<code>expf(x)</code>	<code>__expf(x)</code>
<code>exp10f(x)</code>	<code>__exp10f(x)</code>
<code>powf(x, y)</code>	<code>__powf(x, y)</code>

B.8 Texture Functions

For texture functions, a combination of the texture reference's immutable (compile-time) and mutable (runtime) attributes determine how the texture coordinates are interpreted, what processing occurs during the texture fetch, and the return value delivered by the texture fetch (see Sections 3.2.4.1 and 3.2.4.2).

B.8.1 tex1Dfetch()

```
template<class Type>
Type tex1Dfetch(
    texture<Type, 1, cudaReadModeElementType> texRef,
    int x);

float tex1Dfetch(
    texture<unsigned char, 1, cudaReadModeNormalizedFloat> texRef,
    int x);

float tex1Dfetch(
    texture<signed char, 1, cudaReadModeNormalizedFloat> texRef,
    int x);

float tex1Dfetch(
```

```
texture<unsigned short, 1, cudaReadModeNormalizedFloat> texRef,
int x);

float tex1Dfetch(
    texture<signed short, 1, cudaReadModeNormalizedFloat> texRef,
    int x);
```

fetch the region of linear memory bound to texture reference **texRef** using integer texture coordinate **x**. No texture filtering and addressing modes are supported. For integer types, these functions may optionally promote the integer to single-precision floating point.

Besides the functions shown above, 2-, and 4-tuples are supported; for example:

```
float4 tex1Dfetch(
    texture<uchar4, 1, cudaReadModeNormalizedFloat> texRef,
    int x);
```

fetches the region of linear memory bound to texture reference **texRef** using texture coordinate **x**.

B.8.2 tex1D()

```
template<class Type, enum cudaTextureReadMode readMode>
Type tex1D(texture<Type, 1, readMode> texRef,
    float x);
```

fetches the CUDA array bound to texture reference **texRef** using floating-point texture coordinates **x**.

B.8.3 tex2D()

```
template<class Type, enum cudaTextureReadMode readMode>
Type tex2D(texture<Type, 2, readMode> texRef,
    float x, float y);
```

fetches the CUDA array or the region of linear memory bound to texture reference **texRef** using texture coordinates **x** and **y**.

B.8.4 tex3D()

```
template<class Type, enum cudaTextureReadMode readMode>
Type tex3D(texture<Type, 3, readMode> texRef,
    float x, float y, float z);
```

fetches the CUDA array bound to texture reference **texRef** using texture coordinates **x**, **y**, and **z**.

B.9 Time Function

```
clock_t clock();
```

when executed in device code, returns the value of a per-multiprocessor counter that is incremented every clock cycle. Sampling this counter at the beginning and at the end of a kernel, taking the difference of the two samples, and recording the result per thread provides a measure for each thread of the number of clock cycles

taken by the device to completely execute the thread, but not of the number of clock cycles the device actually spent executing thread instructions. The former number is greater than the latter since threads are time sliced.

B.10 Atomic Functions

An atomic function performs a read-modify-write atomic operation on one 32-bit or 64-bit word residing in global or shared memory. For example, **atomicAdd()** reads a 32-bit word at some address in global or shared memory, adds an integer to it, and writes the result back to the same address. The operation is atomic in the sense that it is guaranteed to be performed without interference from other threads. In other words, no other thread can access this address until the operation is complete.

Atomic operations only work with signed and unsigned integers (with the exception of **atomicExch()**, which is also supported for single-precision floating-point numbers).

Atomic functions can only be used in device functions and are only available for devices of compute capability 1.1 and above.

Atomic functions operating on shared memory and atomic functions operating on 64-bit words are only available for devices of compute capability 1.2 and above.

B.10.1 Arithmetic Functions

B.10.1.1 atomicAdd()

```
int atomicAdd(int* address, int val);
unsigned int atomicAdd(unsigned int* address,
                      unsigned int val);
unsigned long long int atomicAdd(unsigned long long int* address,
                               unsigned long long int val);
```

reads the 32-bit or 64-bit word **old** located at the address **address** in global or shared memory, computes **(old + val)**, and stores the result back to memory at the same address. These three operations are performed in one atomic transaction. The function returns **old**.

64-bit words are only supported for global memory.

B.10.1.2 atomicSub()

```
int atomicSub(int* address, int val);
unsigned int atomicSub(unsigned int* address,
                      unsigned int val);
```

reads the 32-bit word **old** located at the address **address** in global or shared memory, computes **(old - val)**, and stores the result back to memory at the same address. These three operations are performed in one atomic transaction. The function returns **old**.

B.10.1.3 atomicExch()

```
int atomicExch(int* address, int val);
unsigned int atomicExch(unsigned int* address,
```

```

        unsigned int val);
unsigned long long int atomicExch(unsigned long long int* address,
                                unsigned long long int val);
float atomicExch(float* address, float val);

```

reads the 32-bit or 64-bit word **old** located at the address **address** in global or shared memory and stores **val** back to memory at the same address. These two operations are performed in one atomic transaction. The function returns **old**.

64-bit words are only supported for global memory.

B.10.1.4 atomicMin()

```

int atomicMin(int* address, int val);
unsigned int atomicMin(unsigned int* address,
                      unsigned int val);

```

reads the 32-bit word **old** located at the address **address** in global or shared memory, computes the minimum of **old** and **val**, and stores the result back to memory at the same address. These three operations are performed in one atomic transaction. The function returns **old**.

B.10.1.5 atomicMax()

```

int atomicMax(int* address, int val);
unsigned int atomicMax(unsigned int* address,
                      unsigned int val);

```

reads the 32-bit word **old** located at the address **address** in global or shared memory, computes the maximum of **old** and **val**, and stores the result back to memory at the same address. These three operations are performed in one atomic transaction. The function returns **old**.

B.10.1.6 atomicInc()

```

unsigned int atomicInc(unsigned int* address,
                     unsigned int val);

```

reads the 32-bit word **old** located at the address **address** in global or shared memory, computes $((old \geq val) ? 0 : (old+1))$, and stores the result back to memory at the same address. These three operations are performed in one atomic transaction. The function returns **old**.

B.10.1.7 atomicDec()

```

unsigned int atomicDec(unsigned int* address,
                     unsigned int val);

```

reads the 32-bit word **old** located at the address **address** in global or shared memory, computes $((old == 0) \vee (old > val)) ? val : (old-1)$, and stores the result back to memory at the same address. These three operations are performed in one atomic transaction. The function returns **old**.

B.10.1.8 atomicCAS()

```

int atomicCAS(int* address, int compare, int val);
unsigned int atomicCAS(unsigned int* address,
                      unsigned int compare,
                      unsigned int val);
unsigned long long int atomicCAS(unsigned long long int* address,
                                unsigned long long int compare,
                                unsigned long long int val);

```

reads the 32-bit or 64-bit word **old** located at the address **address** in global or shared memory, computes **(old == compare ? val : old)**, and stores the result back to memory at the same address. These three operations are performed in one atomic transaction. The function returns **old** (Compare And Swap).

64-bit words are only supported for global memory.

B.10.2 Bitwise Functions

B.10.2.1 atomicAnd()

```
int atomicAnd(int* address, int val);
unsigned int atomicAnd(unsigned int* address,
                      unsigned int val);
```

reads the 32-bit word **old** located at the address **address** in global or shared memory, computes **(old & val)**, and stores the result back to memory at the same address. These three operations are performed in one atomic transaction. The function returns **old**.

B.10.2.2 atomicOr()

```
int atomicOr(int* address, int val);
unsigned int atomicOr(unsigned int* address,
                    unsigned int val);
```

reads the 32-bit word **old** located at the address **address** in global or shared memory, computes **(old | val)**, and stores the result back to memory at the same address. These three operations are performed in one atomic transaction. The function returns **old**.

B.10.2.3 atomicXor()

```
int atomicXor(int* address, int val);
unsigned int atomicXor(unsigned int* address,
                    unsigned int val);
```

reads the 32-bit word **old** located at the address **address** in global or shared memory, computes **(old ^ val)**, and stores the result back to memory at the same address. These three operations are performed in one atomic transaction. The function returns **old**.

B.11 Warp Vote Functions

Warp vote functions are only supported by devices of compute capability 1.2 and higher (see Section 4.1 for the definition of a warp).

```
int __all(int predicate);
```

evaluates **predicate** for all threads of the warp and returns non-zero if and only if **predicate** evaluates to non-zero for all of them.

```
int __any(int predicate);
```

evaluates **predicate** for all threads of the warp and returns non-zero if and only if **predicate** evaluates to non-zero for any of them.

B.12 Execution Configuration

Any call to a `__global__` function must specify the *execution configuration* for that call. The execution configuration defines the dimension of the grid and blocks that will be used to execute the function on the device, as well as the associated stream (see Section 3.3.9.1 for a description of streams).

When using the driver API, the execution configuration is specified through a series of driver function calls as detailed in Section 3.3.3.

When using the runtime API (Section 3.2), the execution configuration is specified by inserting an expression of the form `<<< Dg, Db, Ns, S >>>` between the function name and the parenthesized argument list, where:

- ❑ **Dg** is of type **dim3** (see Section B.3.2) and specifies the dimension and size of the grid, such that **Dg.x** * **Dg.y** equals the number of blocks being launched; **Dg.z** must be equal to 1;
- ❑ **Db** is of type **dim3** (see Section B.3.2) and specifies the dimension and size of each block, such that **Db.x** * **Db.y** * **Db.z** equals the number of threads per block;
- ❑ **Ns** is of type **size_t** and specifies the number of bytes in shared memory that is dynamically allocated per block for this call in addition to the statically allocated memory; this dynamically allocated memory is used by any of the variables declared as an external array as mentioned in Section B.2.3; **Ns** is an optional argument which defaults to 0;
- ❑ **S** is of type **cudaStream_t** and specifies the associated stream; **S** is an optional argument which defaults to 0.

As an example, a function declared as

```
__global__ void Func(float* parameter);
```

must be called like this:

```
Func<<< Dg, Db, Ns >>>(parameter);
```

The arguments to the execution configuration are evaluated before the actual function arguments and like the function arguments, are currently passed via shared memory to the device.

The function call will fail if **Dg** or **Db** are greater than the maximum sizes allowed for the device as specified in Appendix A.1.1, or if **Ns** is greater than the maximum amount of shared memory available on the device, minus the amount of shared memory required for static allocation, functions arguments, and execution configuration.

Appendix C.

Mathematical Functions

Functions from Section C.1 can be used in both host and device code whereas functions from Section C.2 can only be used in device code.

Note that floating-point functions are overloaded, so that in general, there are three prototypes for a given function **<func-name>**:

- (1) **double <func-name>(double)**, e.g. **double log(double)**
- (2) **float <func-name>(float)**, e.g. **float log(float)**
- (3) **float <func-name>f(float)**, e.g. **float logf(float)**

This means, in particular, that passing a **float** argument always results in a **float** result (variants (2) and (3) above).

C.1 Standard Functions

This section lists all the mathematical standard library functions supported in device code. It also specifies the error bounds of each function when executed on the device. These error bounds also apply when the function is executed on the host in the case where the host does not supply the function. They are generated from extensive but not exhaustive tests, so they are not guaranteed bounds.

C.1.1 Single-Precision Floating-Point Functions

Addition and multiplication are IEEE-compliant, so have a maximum error of 0.5 ulp. However, on the device, the compiler often combines them into a single multiply-add instruction (FMAD), which truncates the intermediate result of the multiplication. This combination can be avoided by using the **__fadd_rn()** and **__fmul_rn()** intrinsic functions (see Section C.2).

The recommended way to round a single-precision floating-point operand to an integer, with the result being a single-precision floating-point number is **rintf()**, not **roundf()**. The reason is that **roundf()** maps to an 8-instruction sequence on the device, whereas **rintf()** maps to a single instruction. **truncf()**, **ceilf()**, and **floorf()** each map to a single instruction as well.

Table C-1. Mathematical Standard Library Functions with Maximum ULP Error

The maximum error is stated as the absolute value of the difference in ulps between a correctly rounded single-precision result and the result returned by the CUDA library function.

Function	Maximum ulp error
x+y	0 (IEEE-754 round-to-nearest-even) (except when merged into an FMAD)
x*y	0 (IEEE-754 round-to-nearest-even) (except when merged into an FMAD)
x/y	2 (full range)
1/x	1 (full range)
1/sqrtf(x) rsqrtf(x)	2 (full range)
sqrtf(x)	3 (full range)
cbrtf(x)	1 (full range)
hypotf(x, y)	3 (full range)
expf(x)	2 (full range)
exp2f(x)	2 (full range)
exp10f(x)	2 (full range)
expm1f(x)	1 (full range)
logf(x)	1 (full range)
log2f(x)	3 (full range)
log10f(x)	3 (full range)
log1pf(x)	2 (full range)
sinf(x)	2 (full range)
cosf(x)	2 (full range)
tanf(x)	4 (full range)
sincosf(x, sptr, cptr)	2 (full range)
asinf(x)	4 (full range)
acosf(x)	3 (full range)
atanf(x)	2 (full range)
atan2f(y, x)	3 (full range)
sinhf(x)	3 (full range)
coshf(x)	2 (full range)
tanhf(x)	2 (full range)
asinhf(x)	3 (full range)
acoshf(x)	4 (full range)
atanhf(x)	3 (full range)
powf(x, y)	8 (full range)
erff(x)	3 (full range)
erfcf(x)	8 (full range)
erfinvf(x)	5 (full range)

Function	Maximum ulp error
<code>erfcinvf(x)</code>	7 (full range)
<code>lgammaf(x)</code>	6 (outside interval -10.001 ... -2.264; larger inside)
<code>tgammaf(x)</code>	11 (full range)
<code>fmaf(x, y, z)</code>	0 (full range)
<code>frexpf(x, exp)</code>	0 (full range)
<code>ldexpf(x, exp)</code>	0 (full range)
<code>scalbnf(x, n)</code>	0 (full range)
<code>scalblnf(x, l)</code>	0 (full range)
<code>logbf(x)</code>	0 (full range)
<code>ilogbf(x)</code>	0 (full range)
<code>fmodf(x, y)</code>	0 (full range)
<code>remainderf(x, y)</code>	0 (full range)
<code>remquof(x, y, iptr)</code>	0 (full range)
<code>modff(x, iptr)</code>	0 (full range)
<code>fdimf(x, y)</code>	0 (full range)
<code>truncf(x)</code>	0 (full range)
<code>roundf(x)</code>	0 (full range)
<code>rintf(x)</code>	0 (full range)
<code>nearbyintf(x)</code>	0 (full range)
<code>ceilf(x)</code>	0 (full range)
<code>floorf(x)</code>	0 (full range)
<code>lrintf(x)</code>	0 (full range)
<code>lroundf(x)</code>	0 (full range)
<code>llrintf(x)</code>	0 (full range)
<code>llroundf(x)</code>	0 (full range)
<code>signbit(x)</code>	N/A
<code>isinf(x)</code>	N/A
<code>isnan(x)</code>	N/A
<code>isfinite(x)</code>	N/A
<code>copysignf(x, y)</code>	N/A
<code>fminf(x, y)</code>	N/A
<code>fmaxf(x, y)</code>	N/A
<code>fabsf(x)</code>	N/A
<code>nanf(cptr)</code>	N/A
<code>nextafterf(x, y)</code>	N/A

C.1.2 Double-Precision Floating-Point Functions

The errors listed below only apply when compiling for devices with native double-precision support. When compiling for devices without such support, such as devices of compute capability 1.2 and lower, the **double** type gets demoted to **float** by default and the double-precision math functions are mapped to their single-precision equivalents.

The recommended way to round a double-precision floating-point operand to an integer, with the result being a double-precision floating-point number is **rint()**, not **round()**. The reason is that **round()** maps to an 8-instruction sequence on the device, whereas **rint()** maps to a single instruction. **trunc()**, **ceil()**, and **floor()** each map to a single instruction as well.

Table C-2. Mathematical Standard Library Functions with Maximum ULP Error

The maximum error is stated as the absolute value of the difference in ulps between a correctly rounded double-precision result and the result returned by the CUDA library function.

Function	Maximum ulp error
x+y	0 (IEEE-754 round-to-nearest-even)
x*y	0 (IEEE-754 round-to-nearest-even)
x/y	0 (IEEE-754 round-to-nearest-even)
1/x	0 (IEEE-754 round-to-nearest-even)
sqrt(x)	0 (IEEE-754 round-to-nearest-even)
rsqrt(x)	1 (full range)
cbrt(x)	1 (full range)
hypot(x, y)	2 (full range)
exp(x)	1 (full range)
exp2(x)	1 (full range)
exp10(x)	1 (full range)
expm1(x)	1 (full range)
log(x)	1 (full range)
log2(x)	1 (full range)
log10(x)	1 (full range)
log1p(x)	1 (full range)
sin(x)	2 (full range)
cos(x)	2 (full range)
tan(x)	2 (full range)
sincos(x, sptr, cptr)	2 (full range)
asin(x)	2 (full range)
acos(x)	2 (full range)
atan(x)	2 (full range)
atan2(y, x)	2 (full range)
sinh(x)	1 (full range)
cosh(x)	1 (full range)
tanh(x)	1 (full range)
asinh(x)	2 (full range)
acosh(x)	2 (full range)
atanh(x)	2 (full range)
pow(x, y)	2 (full range)
erf(x)	2 (full range)

Function	Maximum ulp error
erfc(x)	7 (full range)
erfinv(x)	8 (full range)
erfcinv(x)	8 (full range)
lgamma(x)	4 (outside interval -11.0001 ... -2.2637; larger inside)
tgamma(x)	8 (full range)
fma(x, y, z)	0 (IEEE-754 round-to-nearest-even)
frexp(x, exp)	0 (full range)
ldexp(x, exp)	0 (full range)
scalbn(x, n)	0 (full range)
scalbln(x, l)	0 (full range)
logb(x)	0 (full range)
ilogb(x)	0 (full range)
fmod(x, y)	0 (full range)
remainder(x, y)	0 (full range)
remquo(x, y, iptr)	0 (full range)
modf(x, iptr)	0 (full range)
fdim(x, y)	0 (full range)
trunc(x)	0 (full range)
round(x)	0 (full range)
rint(x)	0 (full range)
nearbyint(x)	0 (full range)
ceil(x)	0 (full range)
floor(x)	0 (full range)
lrint(x)	0 (full range)
lround(x)	0 (full range)
llrint(x)	0 (full range)
llround(x)	0 (full range)
signbit(x)	N/A
isinf(x)	N/A
isnan(x)	N/A
isfinite(x)	N/A
copysign(x, y)	N/A
fmin(x, y)	N/A
fmax(x, y)	N/A
fabs(x)	N/A
nan(cptr)	N/A
nextafter(x, y)	N/A

C.1.3 Integer Functions

Integer **min(x, y)** and **max(x, y)** are supported and map to a single instruction on the device.

C.2 Intrinsic Functions

This section lists the intrinsic functions that are only supported in device code. Among these functions are the less accurate, but faster versions of some of the functions of Section C.1; they have the same name prefixed with `__` (such as `__sinf(x)`).

Functions suffixed with `_rn` operate using the round-to-nearest-even rounding mode.

Functions suffixed with `_rz` operate using the round-towards-zero rounding mode.

Functions suffixed with `_ru` operate using the round-up (to positive infinity) rounding mode.

Functions suffixed with `_rd` operate using the round-down (to negative infinity) rounding mode.

Unlike type conversion functions (such as `__int2float_rn`) that convert from one type to another, type casting functions simply perform a type cast on the argument, leaving the value unchanged. For example, `__int_as_float(0xC0000000)` is equal to `-2`, `__float_as_int(1.0f)` is equal to `0x3f800000`.

C.2.1 Single-Precision Floating-Point Functions

`__fadd_rn()` and `__fmul_rn()` map to addition and multiplication operations that the compiler never merges into FMADs. By contrast, additions and multiplications generated from the `*` and `+` operators will frequently be combined into FMADs.

Both the regular floating-point division and `__fdivdef(x, y)` have the same accuracy, but for $2^{126} < y < 2^{128}$, `__fdivdef(x, y)` delivers a result of zero, whereas the regular division delivers the correct result to within the accuracy stated in Table C-3. Also, for $2^{126} < y < 2^{128}$, if `x` is infinity, `__fdivdef(x, y)` delivers a NaN (as a result of multiplying infinity by zero), while the regular division returns infinity.

`__saturate(x)` returns 0 if `x` is less than 0, 1 if `x` is more than 1, and `x` otherwise.

`__float211_[rn, rz, ru, rd](x)` (respectively `__float2u11_[rn, rz, ru, rd](x)`) converts single-precision floating-point parameter `x` to 64-bit signed (respectively unsigned) integer with specified IEEE-754 rounding modes.

Table C-3. Single-Precision Floating-Point Intrinsic Functions Supported by the CUDA Runtime Library with Respective Error Bounds

Function	Error bounds
<code>__fadd_[rn, rz, ru, rd](x, y)</code>	IEEE-compliant.

<code>__fmul_[rn,rz,ru,rd](x,y)</code>	IEEE-compliant.
<code>__fmaf_[rn,rz,ru,rd](x,y,z)</code>	IEEE-compliant.
<code>__frcp_[rn,rz,ru,rd](x)</code>	IEEE-compliant.
<code>__fsqrt_[rn,rz,ru,rd](x)</code>	IEEE-compliant.
<code>__fdiv_[rn,rz,ru,rd](x,y)</code>	IEEE-compliant.
<code>__fdividef(x,y)</code>	For y in $[2^{-126}, 2^{126}]$, the maximum ulp error is 2.
<code>__expf(x)</code>	The maximum ulp error is 2 + floor(abs(1.16 * x)) .
<code>__exp10f(x)</code>	The maximum ulp error is 2 + floor(abs(2.95 * x)) .
<code>__logf(x)</code>	For x in $[0.5, 2]$, the maximum absolute error is $2^{21.41}$, otherwise, the maximum ulp error is 3.
<code>__log2f(x)</code>	For x in $[0.5, 2]$, the maximum absolute error is 2^{22} , otherwise, the maximum ulp error is 2.
<code>__log10f(x)</code>	For x in $[0.5, 2]$, the maximum absolute error is 2^{24} , otherwise, the maximum ulp error is 3.
<code>__sinf(x)</code>	For x in $[-\pi, \pi]$, the maximum absolute error is $2^{21.41}$, and larger otherwise.
<code>__cosf(x)</code>	For x in $[-\pi, \pi]$, the maximum absolute error is $2^{21.19}$, and larger otherwise.
<code>__sincosf(x, sptr, cptr)</code>	Same as sinf(x) and cosf(x) .
<code>__tanf(x)</code>	Derived from its implementation as __sinf(x) * (1 / __cosf(x)) .
<code>__powf(x, y)</code>	Derived from its implementation as exp2f(y * __log2f(x)) .
<code>__int_as_float(x)</code>	N/A
<code>__float_as_int(x)</code>	N/A
<code>__saturate(x)</code>	N/A
<code>__float2int_[rn,rz,ru,rd](x)</code>	N/A
<code>__float2uint_[rn,rz,ru,rd](x)</code>	N/A
<code>__int2float_[rn,rz,ru,rd](x)</code>	N/A
<code>__uint2float_[rn,rz,ru,rd](x)</code>	N/A
<code>__float2ll_[rn,rz,ru,rd](x)</code>	N/A
<code>__float2ull_[rn,rz,ru,rd](x)</code>	N/A

C.2.2 Double-Precision Floating-Point Functions

`__dadd_rn()` and `__dmul_rn()` map to addition and multiplication operations that the compiler never merges into FMADs. By contrast, additions and multiplications generated from the '*' and '+' operators will frequently be combined into FMADs.

Table C-4. Double-Precision Floating-Point Intrinsic Functions Supported by the CUDA Runtime Library with Respective Error Bounds

Function	Error bounds
<code>__dadd_[rn,rz,ru,rd](x,y)</code>	IEEE-compliant.
<code>__dmul_[rn,rz,ru,rd](x,y)</code>	IEEE-compliant.
<code>__fma_[rn,rz,ru,rd](x,y,z)</code>	IEEE-compliant.
<code>__double2float_[rn,rz](x)</code>	N/A
<code>__double2int_[rn,rz,ru,rd](x)</code>	N/A
<code>__double2uint_[rn,rz,ru,rd](x)</code>	N/A
<code>__double2ll_[rn,rz,ru,rd](x)</code>	N/A
<code>__double2ull_[rn,rz,ru,rd](x)</code>	N/A
<code>__int2double_rn(x)</code>	N/A
<code>__uint2double_rn(x)</code>	N/A
<code>__ll2double_[rn,rz,ru,rd](x)</code>	N/A
<code>__ull2double_[rn,rz,ru,rd](x)</code>	N/A
<code>__double_as_longlong(x)</code>	N/A
<code>__longlong_as_double(x)</code>	N/A
<code>__double2hiint(x)</code>	N/A
<code>__double2loint(x)</code>	N/A
<code>__hioint2double(x, ys)</code>	N/A

C.2.3 Integer Functions

`__[u]mul24(x,y)` computes the product of the 24 least significant bits of the integer parameters **x** and **y** and delivers the 32 least significant bits of the result. The 8 most significant bits of **x** or **y** are ignored.

`__[u]mulhi(x,y)` computes the product of the integer parameters **x** and **y** and delivers the 32 most significant bits of the 64-bit result.

`__[u]mul64hi(x,y)` computes the product of the 64-bit integer parameters **x** and **y** and delivers the 64 most significant bits of the 128-bit result.

`__[u]sad(x,y,z)` (Sum of Absolute Difference) returns the sum of integer parameter **z** and the absolute value of the difference between integer parameters **x** and **y**.

`__clz(x)` returns the number, between 0 and 32 inclusive, of consecutive zero bits starting at the most significant bit (i.e. bit 31) of integer parameter **x**.

`__clzll(x)` returns the number, between 0 and 64 inclusive, of consecutive zero bits starting at the most significant bit (i.e. bit 63) of 64-bit integer parameter **x**.

`__ffs(x)` returns the position of the first (least significant) bit set in integer parameter **x**. The least significant bit is position 1. If **x** is 0, `__ffs()` returns 0. Note that this is identical to the Linux function **ffs**.

__ffsll(x) returns the position of the first (least significant) bit set in 64-bit integer parameter **x**. The least significant bit is position 1. If **x** is 0, **__ffsll()** returns 0. Note that this is identical to the Linux function **ffsll**.

__popc(x) returns the number of bits that are set to 1 in the binary representation of 32-bit integer parameter **x**.

__popcll(x) returns the number of bits that are set to 1 in the binary representation of 64-bit integer parameter **x**.

__brev(x) reverses the bits of 32-bit unsigned integer parameter **x**, i.e. bit N of the result corresponds to bit 31-N of **x**.

__brevll(x) reverses the bits of 64-bit unsigned long long parameter **x**, i.e. bit N of the result corresponds to bit 63-N of **x**.

Appendix D.

C++ Language Constructs

CUDA supports the following C++ language constructs for device code:

- ❑ Polymorphism
- ❑ Default Parameters
- ❑ Operator Overloading
- ❑ Namespaces
- ❑ Function Templates

These C++ constructs are implemented as specified in “The C++ Programming Language” reference. It is valid to use any of these constructs in .cu CUDA files for host, device, and kernel (`__global__`) functions. Any restrictions detailed in previous parts of this programming guide, like the lack of support for recursion, still apply.

The following subsections provide examples of the various constructs.

D.1 Polymorphism

Generally, polymorphism is the ability to define that functions or operators behave differently in different contexts. This is also referred to as function (and operator, see below) overloading.

In practical terms, this means that it is permissible to define two different functions within the same scope (namespace) as long as they have a distinguishable function signature. That means that the two functions either consume a different number of parameters or parameters of different types. When either of the multiple functions gets invoked the compiler resolves to the function’s implementation that matches the function signature.

Because of implicit typecasting, a compiler may encounter multiple potential matches for a function invocation and in that case the matching rules as described in the C++ Language Standard apply. In practice this means that the compiler will pick the closest match in case of multiple potential matches.

Example: The following is valid CUDA code:

```
__device__ void f(float x)
{
    // do something with x
}
```

```

}

__device__ void f(int i)
{
    // do something with i
}

__device__ void f(double x, double y)
{
    // do something with x and y
}

```

D.2 Default Parameters

With support for polymorphism as described in the previous subsection and the function signature matching rules in place it becomes possible to provide support for default values for function parameters.

Example:

```

__device__ void f(float x = 0.0f)
{
    // do something with x
}

```

Kernel or other device functions can now invoke this version of `f` in one of two ways:

```

f();
// or
float x = /* some value */;
f(x);

```

Default parameters can only be given for the last `n` parameters of a function.

D.3 Operator Overloading

Operator overloading allows programmers to define operators for new data-types. Examples of overloadable operators in C++ are: `+`, `-`, `*`, `/`, `+=`, `&`, `[]`, etc.

Example: The following is valid CUDA code, implementing the `+` operation between two `uchar4` vectors:

```

__device__ uchar4 operator+ (const uchar4 & a, const uchar4 & b)
{
    uchar4 r;
    r.x = a.x + b.x;
    ...
    return r;
}

```

This new operator can now be used like this:

```

uchar4 a, b, c;

```

```
a = b = /* some initial value */;
c = a + b;
```

D.4 Namespaces

Namespaces in C++ allow for the creation of a hierarchy of scopes of visibility. All the symbols inside a namespace can be used within this namespaces without additional syntax.

The use of namespaces can be used to solve the problem of name-clashes (two different symbols using identical names), which commonly occurs when using multiple function libraries from different sources.

Example: The following code defines two functions “f()” in two separate namespaces (“nvidia” and “other”):

```
namespace nvidia {
    __device__ void f(float x)
    { /* do something with x */ ;}
}

namespace other {
    __device__ void f(float x)
    { /* do something with x */ ;}
}
```

The functions can now be used anywhere via fully qualified names:

```
nvidia::f(0.5f);
```

All the symbols in a namespace can be imported into another namespace (scope) like this:

```
using namespace nvidia;
f(0.5f);
```

D.5 Function Templates

Function templates are a form of meta-programming that allows writing a generic function in a data-type independent fashion. CUDA supports function templates to the full extent of the C++ standard, including the following concepts:

- ❑ Implicit template parameter decuction.
- ❑ Explicit instantiation.
- ❑ Template specialization.

Example:

```
template <T>
__device__ bool f(T x)
{ return /* some clever code that turns x into a bool here */ }
```

This function will convert x of any data-type to a bool as long as the code in the function’s body can be compiled for the actually type (T) of the variable x.

`f()` can be invoked in two ways:

```
int x = 1;
bool result = f(x);
```

This first type of invocation relies on the compiler's ability to **implicitly deduce** the correct function type for `T`. In this case the compiler would deduce `T` to be `int` and instantiate `f<int>(x)`.

The second type of invoking the template function is via **explicit instantiation** like this:

```
bool result = f<double>(0.5);
```

Function templates may be **specialized**:

```
template <T>
__device__ bool f(T x)
{ return false; }

template <>
__device__ bool
f<int>(T x)
{ return true; }
```

In this case the implementation for `T` representing the `int` type are specialized to return `true`, all other types will be caught by the more general template and return `false`.

The complete set of matching rules (for implicitly deducing template parameters) and matching polymorphous functions apply as specified in the C++ standard.

Appendix E. Texture Fetching

This appendix gives the formula used to compute the value returned by the texture functions of Section B.8 depending on the various attributes of the texture reference (see Section 3.2.4).

The texture bound to the texture reference is represented as an array T of N texels for a one-dimensional texture, $N \times M$ texels for a two-dimensional texture, or $N \times M \times L$ texels for a three-dimensional texture. It is fetched using texture coordinates x , y , and z .

A texture coordinate must fall within T 's valid addressing range before it can be used to address T . The addressing mode specifies how an out-of-range texture coordinate x is remapped to the valid range. If x is non-normalized, only the clamp addressing mode is supported and x is replaced by 0 if $x < 0$ and $N-1$ if $N \leq x$. If x is normalized:

- In clamp addressing mode, x is replaced by 0 if $x < 0$ and $1 - \frac{1}{N}$ if $1 \leq x$,
- In wrap addressing mode, x is replaced by $\text{frac}(x)$, where
$$\text{frac}(x) = x - \text{floor}(x)$$
 and $\text{floor}(x)$ is the largest integer not greater than x .

In the remaining of the appendix, x , y , and z are the non-normalized texture coordinates remapped to T 's valid addressing range. x , y , and z are derived from the normalized texture coordinates $\hat{x}$, $\hat{y}$, and $\hat{z}$ as such: $x = N\hat{x}$, $y = M\hat{y}$, and $z = L\hat{z}$.

E.1 Nearest-Point Sampling

In this filtering mode, the value returned by the texture fetch is

- $\text{tex}(x) = T[i]$ for a one-dimensional texture,
- $\text{tex}(x, y) = T[i, j]$ for a two-dimensional texture,
- $\text{tex}(x, y, z) = T[i, j, k]$ for a three-dimensional texture,

where $i = \text{floor}(x)$, $j = \text{floor}(y)$, and $k = \text{floor}(z)$.

Figure D-1 illustrates nearest-point sampling for a one-dimensional texture with $N = 4$.

For integer textures, the value returned by the texture fetch can be optionally remapped to $[0.0, 1.0]$ (see Section 3.2.4.1).

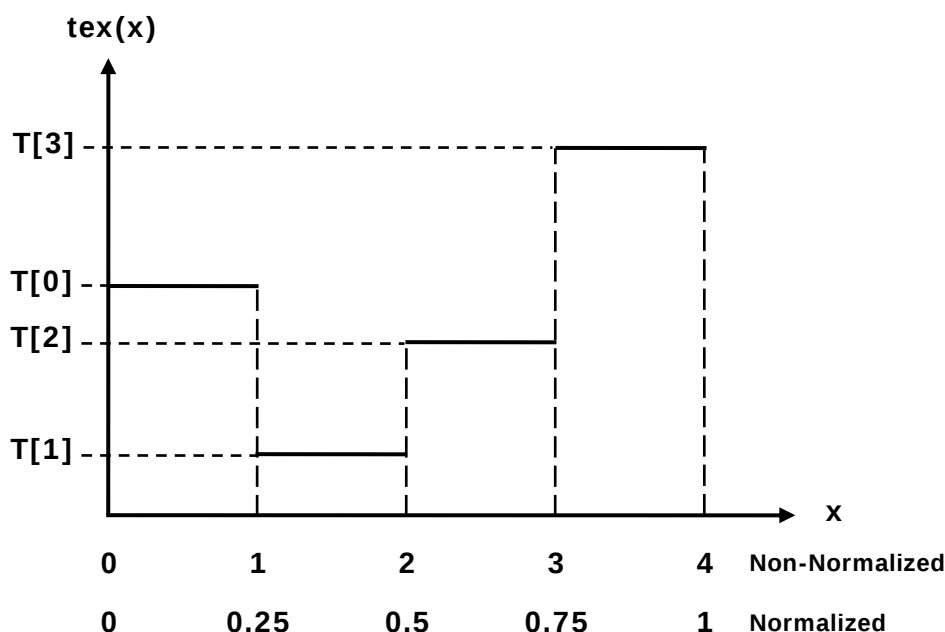


Figure D-1. Nearest-Point Sampling of a One-Dimensional Texture of Four Texels

E.2 Linear Filtering

In this filtering mode, which is only available for floating-point textures, the value returned by the texture fetch is

- $\text{tex}(x) = (1 - \alpha)T[i] + \alpha T[i + 1]$ for a one-dimensional texture,

□ $tex(x, y) = (1 - \alpha)(1 - \beta)T[i, j] + \alpha(1 - \beta)T[i + 1, j] + (1 - \alpha)\beta T[i, j + 1] + \alpha\beta T[i + 1, j + 1]$
for a two-dimensional texture,

□ $tex(x, y, z) =$
 $(1 - \alpha)(1 - \beta)(1 - \gamma)T[i, j, k] + \alpha(1 - \beta)(1 - \gamma)T[i + 1, j, k] +$
 $(1 - \alpha)\beta(1 - \gamma)T[i, j + 1, k] + \alpha\beta(1 - \gamma)T[i + 1, j + 1, k] +$
 $(1 - \alpha)(1 - \beta)\gamma T[i, j, k + 1] + \alpha(1 - \beta)\gamma T[i + 1, j, k + 1] +$
 $(1 - \alpha)\beta\gamma T[i, j + 1, k + 1] + \alpha\beta\gamma T[i + 1, j + 1, k + 1]$

for a three-dimensional texture,

where:

□ $i = \text{floor}(x_B), \alpha = \text{frac}(x_B), x_B = x - 0.5,$

□ $j = \text{floor}(y_B), \beta = \text{frac}(y_B), y_B = y - 0.5,$

□ $k = \text{floor}(z_B), \gamma = \text{frac}(z_B), z_B = z - 0.5.$

α , β , and γ are stored in 9-bit fixed point format with 8 bits of fractional value (so 1.0 is exactly represented).

Figure F-2 illustrates nearest-point sampling for a one-dimensional texture with $N = 4$.

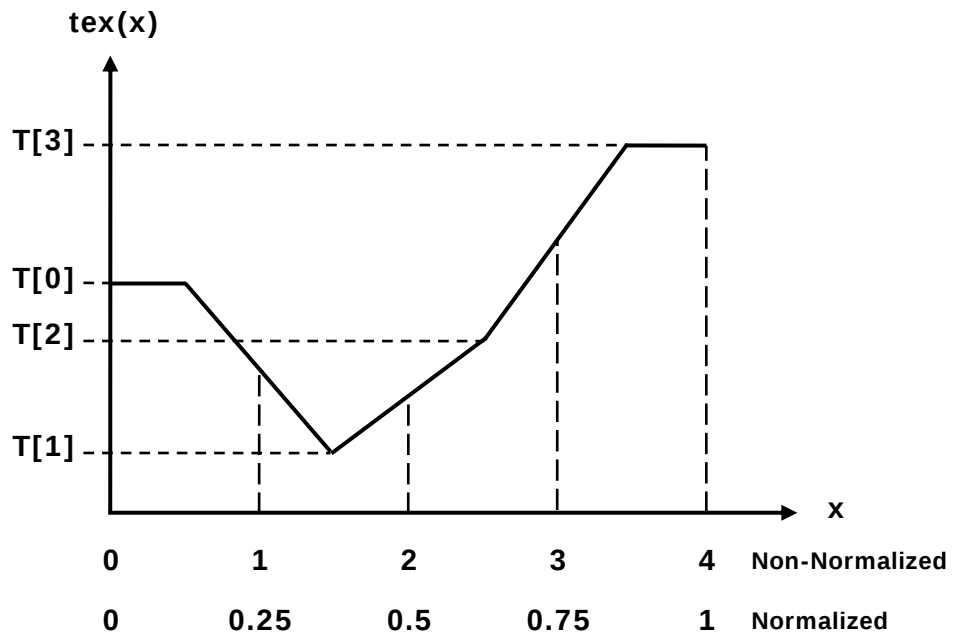


Figure D-2. Linear Filtering of a One-Dimensional Texture of Four Texels in Clamp Addressing Mode

E.3 Table Lookup

A table lookup $TL(x)$ where x spans the interval $[0, R]$ can be implemented as

$$TL(x) = tex\left(\frac{N-1}{R}x + 0.5\right) \text{ in order to ensure that } TL(0) = T[0] \text{ and } TL(R) = T[N-1].$$

Figure F-3 illustrates the use of texture filtering to implement a table lookup with $R = 4$ or $R = 1$ from a one-dimensional texture with $N = 4$.

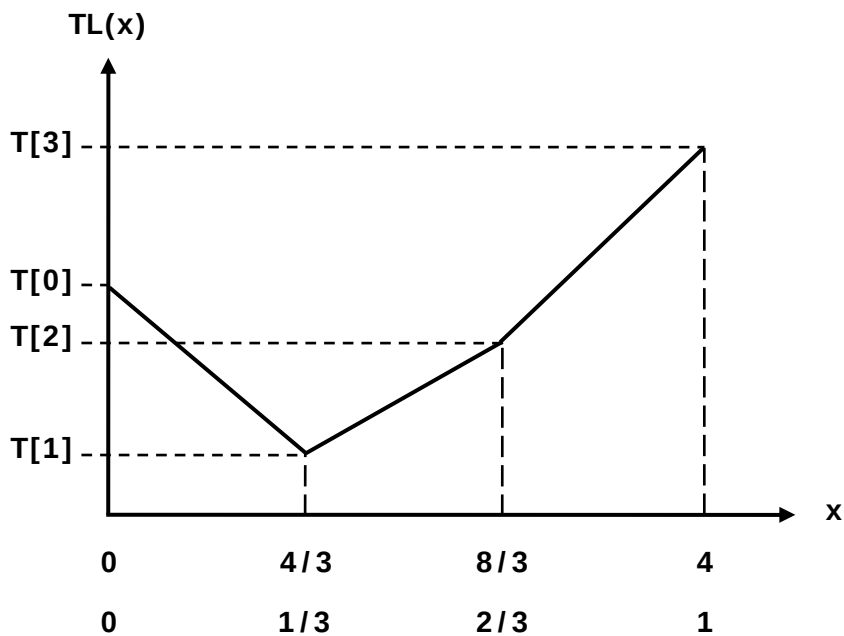


Figure D-3. One-Dimensional Table Lookup Using Linear Filtering

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 or otherwise under any patent or patent rights of NVIDIA Corporation. Specifications mentioned in this publication are subject to change without notice. This publication supersedes and replaces all information previously supplied. NVIDIA Corporation products are not authorized for use as critical components in life support devices or systems without express written approval of NVIDIA Corporation.

Trademarks

NVIDIA, the NVIDIA logo, GeForce, Tesla, and Quadro are trademarks or registered trademarks of NVIDIA Corporation. Other company and product names may be trademarks of the respective companies with which they are associated.

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

Copyright

© 2006-2009 NVIDIA Corporation. All rights reserved.

This work incorporates portions of on an earlier work: Scalable Parallel Programming with CUDA, in ACM Queue, VOL 6, No. 2 (March/April 2008), © ACM, 2008. <http://mags.acm.org/queue/20080304/?u1=texterity>



nvidia.

NVIDIA Corporation
2701 San Tomas Expressway
Santa Clara, CA 95050
www.nvidia.com