



Maxwell Tuning Guide

Release 12.6

NVIDIA Corporation

Nov 14, 2024

Contents

1	NVIDIA Maxwell Compute Architecture	3
2	CUDA Best Practices	5
3	Application Compatibility	7
4	Maxwell Tuning	9
4.1	SMM	9
4.1.1	Occupancy	9
4.1.2	Instruction Scheduling	9
4.1.3	Instruction Latencies	10
4.1.4	Instruction Throughput	10
4.2	Memory Throughput	10
4.2.1	Unified L1/Texture Cache	10
4.3	Shared Memory	11
4.3.1	Shared Memory Capacity	11
4.3.2	Shared Memory Bandwidth	12
4.3.3	Fast Shared Memory Atomics	12
4.4	Dynamic Parallelism	12
5	Revision History	13
6	Notices	15
6.1	Notice	15
6.2	OpenCL	16
6.3	Trademarks	16

Tuning CUDA Applications for Maxwell

The programming guide to tuning CUDA Applications for GPUs based on the NVIDIA Maxwell Architecture.

Chapter 1. NVIDIA Maxwell Compute Architecture

Maxwell is NVIDIA's next-generation architecture for CUDA compute applications. Maxwell retains and extends the same CUDA programming model as in previous NVIDIA architectures such as Fermi and Kepler, and applications that follow the best practices for those architectures should typically see speedups on the Maxwell architecture without any code changes. This guide summarizes the ways that an application can be fine-tuned to gain additional speedups by leveraging Maxwell architectural features.¹

Maxwell introduces an all-new design for the Streaming Multiprocessor (SM) that dramatically improves energy efficiency. Although the Kepler SMX design was extremely efficient for its generation, through its development, NVIDIA's GPU architects saw an opportunity for another big leap forward in architectural efficiency; the Maxwell SM is the realization of that vision. Improvements to control logic partitioning, workload balancing, clock-gating granularity, compiler-based scheduling, number of instructions issued per clock cycle, and many other enhancements allow the Maxwell SM (also called *SMM*) to far exceed Kepler SMX efficiency.

The first Maxwell-based GPU is codenamed *GM107* and is designed for use in power-limited environments like notebooks and small form factor (SFF) PCs. *GM107* is described in a whitepaper entitled [NVIDIA GeForce GTX 750 Ti: Featuring First-Generation Maxwell GPU Technology, Designed for Extreme Performance per Watt](#).²

The first GPU using the second-generation Maxwell architecture is codenamed *GM204*. Second-generation Maxwell GPUs retain the power efficiency of the earlier generation while delivering significantly higher performance. *GM204* is described in a whitepaper entitled [NVIDIA GeForce GTX 980: Featuring Maxwell, The Most Advanced GPU Ever Made](#).

Compute programming features of *GM204* are similar to those of *GM107*, except where explicitly noted in this guide. For details on the programming features discussed in this guide, please refer to the [CUDA C++ Programming Guide](#).

¹ Throughout this guide, *Fermi* refers to devices of compute capability 2.x, *Kepler* refers to devices of compute capability 3.x, and *Maxwell* refers to devices of compute capability 5.x.

² The features of *GM108* are similar to those of *GM107*.

Chapter 2. CUDA Best Practices

The performance guidelines and best practices described in the [CUDA C++ Programming Guide](#) and the [CUDA C++ Best Practices Guide](#) apply to all CUDA-capable GPU architectures. Programmers must primarily focus on following those recommendations to achieve the best performance.

The high-priority recommendations from those guides are as follows:

- ▶ Find ways to parallelize sequential code,
- ▶ Minimize data transfers between the host and the device,
- ▶ Adjust kernel launch configuration to maximize device utilization,
- ▶ Ensure global memory accesses are coalesced,
- ▶ Minimize redundant accesses to global memory whenever possible,
- ▶ Avoid long sequences of diverged execution by threads within the same warp.

Chapter 3. Application Compatibility

Before addressing specific performance tuning issues covered in this guide, refer to the [Maxwell Compatibility Guide for CUDA Applications](#) to ensure that your application is compiled in a way that is compatible with Maxwell.

Chapter 4. Maxwell Tuning

4.1. SMM

The Maxwell Streaming Multiprocessor, SMM, is similar in many respects to the Kepler architecture's SMX. The key enhancements of SMM over SMX are geared toward improving efficiency without requiring significant increases in available parallelism per SM from the application.

4.1.1. Occupancy

The maximum number of concurrent warps per SMM remains the same as in SMX (i.e., 64), and [factors influencing warp occupancy](#) remain similar or improved over SMX:

- ▶ The register file size (64k 32-bit registers) is the same as that of SMX.
- ▶ The maximum registers per thread, 255, matches that of Kepler GK110. As with Kepler, experimentation should be used to determine the optimum balance of register spilling vs. occupancy, however.
- ▶ The maximum number of thread blocks per SM has been increased from 16 to 32. This should result in an automatic occupancy improvement for kernels with small thread blocks of 64 or fewer threads (shared memory and register file resource requirements permitting). Such kernels would have tended to under-utilize SMX, but less so SMM.
- ▶ Shared memory capacity is increased (see [Shared Memory Capacity](#)).

As such, developers can expect similar or improved occupancy on SMM without changes to their application. At the same time, warp occupancy requirements (i.e., available parallelism) for maximum device utilization are similar to or less than those of SMX (see [Instruction Latencies](#)).

4.1.2. Instruction Scheduling

The number of CUDA Cores per SM has been reduced to a power of two, however with Maxwell's improved execution efficiency, performance per SM is usually within 10% of Kepler performance, and the improved area efficiency of SMM means CUDA Cores per GPU will be substantially higher vs. comparable Fermi or Kepler chips. SMM retains the same number of instruction issue slots per clock and reduces arithmetic latencies compared to the Kepler design.

As with SMX, each SMM has four warp schedulers. Unlike SMX, however, all SMM core functional units are assigned to a particular scheduler, with no shared units. Along with the selection of a power-of-two

number of CUDA Cores per SM, which simplifies scheduling and reduces stall cycles, this partitioning of SM computational resources in SMM is a major component of the streamlined efficiency of SMM.

The power-of-two number of CUDA Cores per partition simplifies scheduling, as each of SMM's warp schedulers issue to a dedicated set of CUDA Cores equal to the warp width. Each warp scheduler still has the flexibility to dual-issue (such as issuing a math operation to a CUDA Core in the same cycle as a memory operation to a load/store unit), but single-issue is now sufficient to fully utilize all CUDA Cores.

4.1.3. Instruction Latencies

Another major improvement of SMM is that dependent math latencies have been significantly reduced; a consequence of this is a further reduction of stall cycles, as the available warp-level parallelism (i.e., occupancy) on SMM should be equal to or greater than that of SMX (see [Occupancy](#)), while at the same time each math operation takes *less* time to complete, improving utilization and throughput.

4.1.4. Instruction Throughput

The most significant changes to peak instruction throughputs in SMM are as follows:

- ▶ The change in [number of CUDA Cores per SM](#) brings with it a corresponding change in peak single-precision floating point operations per clock per SM. However, since the number of SMs is typically increased, the result is an increase in aggregate peak throughput; furthermore, the scheduling and latency improvements also discussed above make this peak easier to approach.
- ▶ The throughput of many integer operations including multiply, logical operations and shift is improved. In addition, there are now specialized integer instructions that can accelerate pointer arithmetic. These instructions are most efficient when data structures are a power of two in size.

Note: As was already the recommended best practice, signed arithmetic should be preferred over unsigned arithmetic wherever possible for best throughput on SMM. The C language standard places more restrictions on overflow behavior for unsigned math, limiting compiler optimization opportunities.

4.2. Memory Throughput

4.2.1. Unified L1/Texture Cache

Maxwell combines the functionality of the L1 and texture caches into a single unit.

As with Kepler, global loads in Maxwell are cached in L2 only, unless using the *LDG* read-only data cache mechanism introduced in Kepler.

In a manner similar to Kepler GK1 10B, GM204 retains this behavior by default but also allows applications to opt-in to caching of global loads in its unified L1/Texture cache. The opt-in mechanism is the same as with GK1 10B: pass the `-Xptxas -d1cm=ca` flag to `nvcc` at compile time.

Local loads also are cached in L2 only, which could increase the cost of register spilling if L1 local load hit rates were high with Kepler. The balance of occupancy versus spilling should therefore be reevaluated to ensure best performance. Especially given the improvements to arithmetic latencies, code built for Maxwell may benefit from somewhat lower occupancy (due to increased registers per thread) in exchange for lower spilling.

The unified L1/texture cache acts as a coalescing buffer for memory accesses, gathering up the data requested by the threads of a warp prior to delivery of that data to the warp. This function previously was served by the separate L1 cache in Fermi and Kepler.

Two new device attributes were added in CUDA Toolkit 6.0: `globalL1CacheSupported` and `localL1CacheSupported`. Developers who wish to have separately-tuned paths for various architecture generations can use these fields to simplify the path selection process.

Note: Enabling caching of globals in GM204 can affect occupancy. If per-thread-block SM resource usage would result in zero occupancy with caching enabled, the CUDA driver will override the caching selection to allow the kernel launch to succeed. This situation is reported by the profiler.

4.3. Shared Memory

4.3.1. Shared Memory Capacity

With Fermi and Kepler, shared memory and the L1 cache shared the same on-chip storage. Maxwell, by contrast, provides dedicated space to the shared memory of each SMM, since the functionality of the L1 and texture caches have been merged in SMM. This increases the shared memory space available per SMM as compared to SMX: GM107 provides 64 KB shared memory per SMM, and GM204 further increases this to 96 KB shared memory per SMM.

This presents several benefits to application developers:

- ▶ Algorithms with significant shared memory capacity requirements (e.g., radix sort) see an automatic 33% to 100% boost in capacity per SM on top of the aggregate boost from higher SM count.
- ▶ Applications no longer need to select a preference of the L1/shared split for optimal performance. For purposes of backward compatibility with Fermi and Kepler, applications may optionally continue to specify such a preference, but the preference will be ignored on Maxwell, with the full 64 KB per SMM always going to shared memory.

Note: While the per-SM shared memory capacity is increased in SMM, the per-thread-block limit remains 48 KB. For maximum flexibility on possible future GPUs, NVIDIA recommends that applications use at most 32 KB of shared memory in any one thread block, which would for example allow at least two such thread blocks to fit per SMM.

4.3.2. Shared Memory Bandwidth

Kepler SMX introduced an optional 8-byte shared memory banking mode, which had the potential to increase shared memory bandwidth per SM over Fermi for shared memory accesses of 8 or 16 bytes. However, applications could only benefit from this when storing these larger elements in shared memory (i.e., integers and fp32 values saw no benefit), and only when the developer explicitly opted into the 8-byte bank mode via the API.

To simplify this, Maxwell returns to the Fermi style of shared memory banking, where banks are always four bytes wide. Aggregate shared memory bandwidth across the chip remains comparable to that of corresponding Kepler chips, given increased SM count. In this way, all applications using shared memory can now benefit from the higher bandwidth, even when storing only four-byte items into shared memory and without specifying any particular preference via the API.

4.3.3. Fast Shared Memory Atomics

Kepler introduced a dramatically higher throughput for atomic operations to *global* memory as compared to Fermi. However, atomic operations to *shared* memory remained essentially unchanged: both architectures implemented shared memory atomics using a lock/update/unlock pattern that could be expensive in the case of high contention for updates to particular locations in shared memory.

Maxwell improves upon this by implementing native shared memory atomic operations for 32-bit integers and native shared memory 32-bit and 64-bit compare-and-swap (CAS), which can be used to implement other atomic functions with reduced overhead compared to the Fermi and Kepler methods.

Note: Refer to the [CUDA C++ Programming Guide](#) for an example implementation of an fp64 atomicAdd() using atomicCAS().

4.4. Dynamic Parallelism

GK110 introduced a new architectural feature called Dynamic Parallelism, which allows the GPU to create additional work for itself. A programming model enhancement leveraging this feature was introduced in CUDA 5.0 to enable kernels running on GK110 to launch additional kernels onto the same GPU.

SMM brings Dynamic Parallelism into the mainstream by supporting it across the product line, even in lower-power chips such as GM107. This will benefit developers, as it means that applications will no longer need special-case algorithm implementations for high-end GPUs that differ from those usable in more power-constrained environments.

Chapter 5. Revision History

Version 1.0

- ▶ Initial Public Release

Version 1.1

- ▶ Updated for second-generation Maxwell (compute capability 5.2).

Version 1.2

- ▶ Updated references to the CUDA C++ Programming Guide and CUDA C++ Best Practices Guide.

Chapter 6. Notices

6.1. Notice

This document is provided for information purposes only and shall not be regarded as a warranty of a certain functionality, condition, or quality of a product. NVIDIA Corporation (“NVIDIA”) makes no representations or warranties, expressed or implied, as to the accuracy or completeness of the information contained in this document and assumes no responsibility for any errors contained herein. NVIDIA shall have no liability for the consequences or use of such information or for any infringement of patents or other rights of third parties that may result from its use. This document is not a commitment to develop, release, or deliver any Material (defined below), code, or functionality.

NVIDIA reserves the right to make corrections, modifications, enhancements, improvements, and any other changes to this document, at any time without notice.

Customer should obtain the latest relevant information before placing orders and should verify that such information is current and complete.

NVIDIA products are sold subject to the NVIDIA standard terms and conditions of sale supplied at the time of order acknowledgement, unless otherwise agreed in an individual sales agreement signed by authorized representatives of NVIDIA and customer (“Terms of Sale”). NVIDIA hereby expressly objects to applying any customer general terms and conditions with regards to the purchase of the NVIDIA product referenced in this document. No contractual obligations are formed either directly or indirectly by this document.

NVIDIA products are not designed, authorized, or warranted to be suitable for use in medical, military, aircraft, space, or life support equipment, nor in applications where failure or malfunction of the NVIDIA product can reasonably be expected to result in personal injury, death, or property or environmental damage. NVIDIA accepts no liability for inclusion and/or use of NVIDIA products in such equipment or applications and therefore such inclusion and/or use is at customer’s own risk.

NVIDIA makes no representation or warranty that products based on this document will be suitable for any specified use. Testing of all parameters of each product is not necessarily performed by NVIDIA. It is customer’s sole responsibility to evaluate and determine the applicability of any information contained in this document, ensure the product is suitable and fit for the application planned by customer, and perform the necessary testing for the application in order to avoid a default of the application or the product. Weaknesses in customer’s product designs may affect the quality and reliability of the NVIDIA product and may result in additional or different conditions and/or requirements beyond those contained in this document. NVIDIA accepts no liability related to any default, damage, costs, or problem which may be based on or attributable to: (i) the use of the NVIDIA product in any manner that is contrary to this document or (ii) customer product designs.

No license, either expressed or implied, is granted under any NVIDIA patent right, copyright, or other NVIDIA intellectual property right under this document. Information published by NVIDIA regarding third-party products or services does not constitute a license from NVIDIA to use such products or

services or a warranty or endorsement thereof. Use of such information may require a license from a third party under the patents or other intellectual property rights of the third party, or a license from NVIDIA under the patents or other intellectual property rights of NVIDIA.

Reproduction of information in this document is permissible only if approved in advance by NVIDIA in writing, reproduced without alteration and in full compliance with all applicable export laws and regulations, and accompanied by all associated conditions, limitations, and notices.

THIS DOCUMENT AND ALL NVIDIA DESIGN SPECIFICATIONS, REFERENCE BOARDS, FILES, DRAWINGS, DIAGNOSTICS, LISTS, AND OTHER DOCUMENTS (TOGETHER AND SEPARATELY, "MATERIALS") ARE BEING PROVIDED "AS IS." NVIDIA MAKES NO WARRANTIES, EXPRESSED, IMPLIED, STATUTORY, OR OTHERWISE WITH RESPECT TO THE MATERIALS, AND EXPRESSLY DISCLAIMS ALL IMPLIED WARRANTIES OF NONINFRINGEMENT, MERCHANTABILITY, AND FITNESS FOR A PARTICULAR PURPOSE. TO THE EXTENT NOT PROHIBITED BY LAW, IN NO EVENT WILL NVIDIA BE LIABLE FOR ANY DAMAGES, INCLUDING WITHOUT LIMITATION ANY DIRECT, INDIRECT, SPECIAL, INCIDENTAL, PUNITIVE, OR CONSEQUENTIAL DAMAGES, HOWEVER CAUSED AND REGARDLESS OF THE THEORY OF LIABILITY, ARISING OUT OF ANY USE OF THIS DOCUMENT, EVEN IF NVIDIA HAS BEEN ADVISED OF THE POSSIBILITY OF SUCH DAMAGES. Notwithstanding any damages that customer might incur for any reason whatsoever, NVIDIA's aggregate and cumulative liability towards customer for the products described herein shall be limited in accordance with the Terms of Sale for the product.

6.2. OpenCL

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

6.3. Trademarks

NVIDIA and the NVIDIA logo are trademarks or registered trademarks of NVIDIA Corporation in the U.S. and other countries. Other company and product names may be trademarks of the respective companies with which they are associated.

Copyright

©2014-2024, NVIDIA Corporation & affiliates. All rights reserved