



CUDA BINARY UTILITIES

DA-06762-001_v5.5 | May 2013

Application Note



TABLE OF CONTENTS

Chapter 1. Overview.....	1
1.1 What is a CUDA Binary?	1
1.2 Differences between cuobjdump and nvdiasm	1
Chapter 2. cuobjdump.....	3
2.1 Usage	3
2.2 Command-line Options	5
Chapter 3. nvdiasm.....	6
3.1 Usage	6
3.2 Command-line Options	8
Chapter 4. Instruction Set Reference.....	9
4.1 GT200 Instruction Set	9
4.2 Fermi Instruction Set	11
4.3 Kepler Instruction Set	15

LIST OF FIGURES

Figure 1 Control Flow Graph.....	7
----------------------------------	---

LIST OF TABLES

Table 1	Comparison of cuobjdump and nvdiasm.....	1
Table 2	cuobjdump Command-line Options.....	5
Table 3	nvdiasm Command-line Options.....	8
Table 4	GT200 Instruction Set.....	9
Table 5	Fermi Instruction Set.....	12
Table 6	Kepler Instruction Set.....	15

Chapter 1.

OVERVIEW

This document introduces **cuobjdump** and **nvdiasm**, two CUDA binary tools for Linux, Windows and Mac OS.

1.1 What is a CUDA Binary?

A CUDA binary (also referred to as cubin) file is an ELF-formatted file which consists of CUDA executable code sections as well as other sections containing symbols, relocators, debug info, etc. By default, the CUDA compiler driver **nvcc** embeds cubin files into the host executable file. But they can also be generated separately by using the "**-cubin**" option of **nvcc**. cubin files are loaded at run time by the CUDA driver API.



For more details on cubin files or the CUDA compilation trajectory, refer to [NVIDIA CUDA Compiler Driver NVCC](#).

1.2 Differences between **cuobjdump** and **nvdiasm**

CUDA provides two binary utilities for examining and disassembling cubin files and host executables: **cuobjdump** and **nvdiasm**. Basically, **cuobjdump** accepts both cubin files and host binaries while **nvdiasm** only accepts cubin files; but **nvdiasm** provides richer output options.

Here's a quick comparison of the two tools:

Table 1 Comparison of **cuobjdump** and **nvdiasm**

	cuobjdump	nvdiasm
Disassemble cubin	Yes	Yes

	cuobjdump	nvdiasm
Extract ptx and extract and disassemble cubin from the following input files: ▶ Host binaries ▶ Executables ▶ Object files ▶ Static libraries ▶ External fatbinary files	Yes	No
Control flow analysis and output	No	Yes
Advanced display options	No	Yes



Currently, cuobjdump accepts older fatbinary format generated using nvcc 3.2 or earlier. This support might be deprecated in next release.

Chapter 2.

CUOBJDUMP

cuobjdump extracts information from CUDA binary files (both standalone and those embedded in host binaries) and presents them in human readable format. The output of **cuobjdump** includes CUDA assembly code for each kernel, CUDA ELF section headers, string tables, relocators and other CUDA specific sections. It also extracts embedded ptx text from host binaries.

For a list of CUDA assembly instruction set of each GPU architecture, see [Instruction Set Reference](#).

2.1 Usage

cuobjdump accepts a single input file each time it's run. The basic usage is as following:

```
cuobjdump [options] <file>
```

To disassemble a standalone cubin or cubins embedded in a host executable and show CUDA assembly of the kernels, use the following command:

```
cuobjdump -sass <input file>
```

To dump cuda elf sections in human readable format from a cubin file, use the following command:

```
cuobjdump -elf <cubin file>
```

To extract ptx text from a host binary, use the following command:

```
cuobjdump -ptx <host binary>
```

Here's a sample output of **cuobjdump**:

```
$ cuobjdump -ptx -sass a.out
.....
        code for sm_20
Fatbin elf code:
=====
arch = sm_20
```

```

code version = [1,7]
producer = cuda
host = linux
compile_size = 64bit
identifier = main.cu

        code for sm_20
                Function : _Z4addXPii
.headerflags      @"EF_CUDA_SM20 EF_CUDA_PTX_SM(EF_CUDA_SM20)"
/*0000*/          MOV R1, c[0x1][0x100];          /* 0x2800440400005de4 */
/*0008*/          MOV R0, c[0x0][0x28];          /* 0x28004000a0001de4 */
/*0010*/          MOV R4, c[0x0][0x20];          /* 0x2800400080011de4 */
/*0018*/          MOV R5, c[0x0][0x24];          /* 0x2800400090015de4 */
/*0020*/          ISETP.LT.AND P0, PT, R0, 0x1, PT; /* 0x188ec0000401dc23 */
/*0028*/          @P0 EXIT ;                      /* 0x80000000000001e7 */
/*0030*/          LDU.E R0, [R4];                /* 0x8c00000000401c85 */
/*0038*/          MOV R2, RZ;                    /* 0x28000000fc009de4 */
/*0040*/          IADD R2, R2, 0x1;              /* 0x4800c00004209c03 */
/*0048*/          MOV R3, R0;                    /* 0x280000000000dde4 */
/*0050*/          IADD R0, R0, 0x1;              /* 0x4800c00004001c03 */
/*0058*/          ISETP.LT.AND P0, PT, R0, 0x1, PT; /* 0x188ec0000401dc23 */
/*0060*/          @P0 BRA 0x40;                  /* 0x4003ffff600001e7 */
/*0068*/          IADD R0, R3, 0x1;              /* 0x4800c00004301c03 */
/*0070*/          ST.E [R4], R0;                 /* 0x9400000000401c85 */
/*0078*/          EXIT ;                        /* 0x80000000000001de7 */
                .....

.....

.version 3.2
.target sm_30
.address_size 64

.file 1 "/home/cuda/test.cu", 1363124876, 83

.visible .entry _Z4addXPii(
.param .u64 _Z4addXPii_param_0,
.param .u32 _Z4addXPii_param_1
)
{
.reg .pred %p<3>;
.reg .s32 %r<11>;
.reg .s64 %rd<3>;

ld.param.u64 %rd2, [_Z4addXPii_param_0];
ld.param.u32 %r6, [_Z4addXPii_param_1];
cvta.to.global.u64 %rd1, %rd2;
.loc 1 4 1
setp.lt.s32 %p1, %r6, 1;
@%p1 bra BB0_4;

.loc 1 5 1
ld.global.u32 %r9, [%rd1];
mov.u32 %r10, 0;

BB0_2:
.loc 1 5 1
mov.u32 %r2, %r9;
add.s32 %r9, %r2, 1;
.loc 1 4 92
add.s32 %r10, %r10, 1;
.loc 1 4 1
setp.lt.s32 %p2, %r10, %r6;
@%p2 bra BB0_2;

.loc 1 5 1

```

```

add.s32 %r8, %r2, 1;
st.global.u32 [%rd1], %r8;

BB0_4:
.loc 1 7 2
ret;
}

```

As is shown in the output, the `a.out` host binary contains cubin for `sm_20` and ptx code for `sm_30`.

2.2 Command-line Options

Table 2 contains supported command-line options of `cuobjdump`, along with a description of what each option does. Each option has a long name and a short name, which can be used interchangeably.

Table 2 `cuobjdump` Command-line Options

Option (long)	Option (short)	Description
<code>--all-fatbin</code>	<code>-all</code>	Dump all fatbin sections. By default will only dump contents of executable fatbin (if exists), else relocatable fatbin if no executable fatbin.
<code>--dump-elf</code>	<code>-elf</code>	Dump ELF Object sections.
<code>--dump-elf-symbols</code>	<code>-symbols</code>	Dump ELF symbol names.
<code>--dump-ptx</code>	<code>-ptx</code>	Dump PTX for all listed device functions.
<code>--dump-sass</code>	<code>-sass</code>	Dump CUDA assembly for a single cubin file or all cubin files embedded in the binary.
<code>--file <filename>,...</code>	<code>-f</code>	Specify names of source files whose fat binary structures must be dumped. Source files may be specified by the full path by which they were compiled using <code>nvcc</code> , or file name only (omitting the directory part), or file base name only (omitting directory and the <code>'.cu'</code> file extension). This option will be deprecated in next release as it is only applicable to fatbinary generated using <code>nvcc</code> 3.2 or earlier.
<code>--function <function name>,...</code>	<code>-fun</code>	Specify names of device functions whose fat binary structures must be dumped.
<code>--help</code>	<code>-h</code>	Print this help information on this tool.
<code>--options-file <file>,...</code>	<code>-optf</code>	Include command line options from specified file.
<code>--sort-functions</code>	<code>-sort</code>	Sort functions when dumping sass.
<code>--version</code>	<code>-V</code>	Print version information on this tool.

Chapter 3.

NVDISASM

nvdiasm extracts information from standalone cubin files and presents them in human readable format. The output of **nvdiasm** includes CUDA assembly code for each kernel, listing of ELF data sections and other CUDA specific sections. Output style and options are controlled through **nvdiasm** command-line options. **nvdiasm** also does control flow analysis to annotate jump/branch targets and makes the output easier to read.



nvdiasm requires complete relocation information to do control flow analysis. If this information is missing from the CUDA binary, either use the **nvdiasm** option **"-ndf"** to turn off control flow analysis, or use the **ptxas** and **nvlink** option **"-preserve-relocs"** to re-generate the cubin file.

For a list of CUDA assembly instruction set of each GPU architecture, see [Instruction Set Reference](#).

3.1 Usage

nvdiasm accepts a single input file each time it's run. The basic usage is as following:

```
nvdiasm [options] <input cubin file>
```

To get the control flow graph of a kernel, use the following:

```
nvdiasm -cfg <input cubin file>
```

Here's a sample output of **nvdiasm**:

```
.headerflags    @"EF_CUDA_TEXMODE_UNIFIED EF_CUDA_64BIT_ADDRESS
EF_CUDA_SM30 EF_CUDA_PTX_SM(EF_CUDA_SM30) "

//----- .nv.info -----
      .section   .nv.info,"",@"SHT_CUDA_INFO "
      .align    4

.....
```

```
//----- .text._Z4addXPii -----
.section .text._Z4addXPii,"ax",@progbits
.sectioninfo @"SHI_REGISTERS=11 "
.align 4
.global _Z4addXPii
.type _Z4addXPii,@function
.size _Z4addXPii,(.L_19 - _Z4addXPii)
.other _Z4addXPii,@"STO_CUDA_ENTRY STV_DEFAULT "
_Z4addXPii:
.text._Z4addXPii:
/*0008*/ MOV R1, c[0x0][0x44];
/*0010*/ ISUB R1, R1, 0x8;
/*0018*/ MOV R0, c[0x0][0x148];
/*0020*/ IADD R6.CC, R1, c[0x0][0x24];
/*0028*/ ISETP.LT.AND P0, PT, R0, 0x1, PT;
/*0030*/ MOV R8, c[0x0][0x140];
/*0038*/ MOV R9, c[0x0][0x144];
/*0048*/ IADD.X R7, RZ, RZ;
/*0050*/ ISUB R10, R6, c[0x0][0x24];
/*0058*/ @P0 BRA `(.L_2);
/*0060*/ LD.E R0, [R8];
/*0068*/ MOV R2, RZ;
/*0070*/ NOP;
/*0078*/ NOP;

.L_3:
/*0088*/ IADD R2, R2, 0x1;
/*0090*/ MOV R3, R0;
/*0098*/ IADD R0, R0, 0x1;
/*00a0*/ ISETP.LT.AND P0, PT, R2, c[0x0][0x148], PT;
/*00a8*/ @P0 BRA `(.L_3);
/*00b0*/ IADD R0, R3, 0x1;
/*00b8*/ ST.E [R8], R0;

.L_2:
/*00c8*/ S2R R0, SR_TID.X;
/*00d0*/ ISETP.NE.AND P0, PT, R0, RZ, PT;
/*00d8*/ @P0 EXIT ;
/*00e0*/ LD.E R0, [R8];
/*00e8*/ MOV R4, c[0x0][0xf0];
/*00f0*/ MOV R5, c[0x0][0xf4];
/*00f8*/ STL [R10], R0;
/*0108*/ JCAL `(vprintf);
/*0110*/ EXIT ;

.L_4:
/*0118*/ BRA `(.L_4);

.L_19:

//----- SYMBOLS -----

.type vprintf,@function
```

nvdisasm is capable of generating control flow of CUDA assembly in the format of DOT graph description language. The output of the control flow from **nvdisasm** can be directly imported to a DOT graph visualization tool such as [Graphviz](#).

Here's how you can generate a PNG image (**cfg.png**) of the control flow of the above cubin (**a.cubin**) with **nvdisasm** and Graphviz:

```
nvdisasm -cfg a.cubin | dot -ocfg.png -Tpng
```

Here's the generated graph:

Figure 1 Control Flow Graph

3.2 Command-line Options

Table 3 contains the supported command-line options of **nvdiasm**, along with a description of what each option does. Each option has a long name and a short name, which can be used interchangeably.

Table 3 nvdiasm Command-line Options

Option (long)	Option (short)	Description
<code>--base-address <value></code>	<code>-base</code>	Specify the logical base address of the image to disassemble. This option is only valid when disassembling a raw instruction binary (see option '--binary'), and is ignored when disassembling an Elf file. Default value: 0.
<code>--binary <SMxy></code>	<code>-b</code>	When this option is specified, the input file is assumed to contain a raw instruction binary, that is, a sequence of binary instruction encodings as they occur in instruction memory. The value of this option must be the asserted architecture of the raw binary. Allowed values for this option: 'SM10', 'SM20', 'SM30', 'SM35'.
<code>--help</code>	<code>-h</code>	Print this help information on this tool.
<code>--no-dataflow</code>	<code>-ndf</code>	Disable dataflow analyzer after disassembly. Dataflow analysis is normally enabled to perform branch stack analysis and annotate all instructions that jump via the GPU branch stack with inferred branch target labels. However, it may occasionally fail when certain restrictions on the input nvelf/cubin are not met.
<code>--options-file <file>,...</code>	<code>-optf</code>	Include command line options from specified file.
<code>--output-control-flow-graph</code>	<code>-cfg</code>	When specified, output the control flow graph in a format consumable by graphviz tools (such as dot).
<code>--print-code</code>	<code>-c</code>	Only print code sections.
<code>--print-instruction-encoding</code>	<code>-hex</code>	When specified, print the encoding bytes after each disassembled operation.
<code>--print-raw</code>	<code>-raw</code>	Print the disassembly without any attempt to beautify it.
<code>--separate-functions</code>	<code>-sf</code>	Separate the code corresponding with function symbols by some new lines to let them stand out in the printed disassembly.
<code>--version</code>	<code>-v</code>	Print version information on this tool.

Chapter 4.

INSTRUCTION SET REFERENCE

This is an instruction set reference of NVIDIA® GPU architectures GT200, Fermi and Kepler.

4.1 GT200 Instruction Set

The GT200 architecture (Compute Capability 1.x) has the following instruction set format:

```
(instruction) (destination) (source1), (source2) ...
```

Valid destination and source locations include:

- ▶ RX for registers
- ▶ AX for address registers
- ▶ SRX for special system-controlled registers
- ▶ CX for condition registers
- ▶ global14 r[X] for global memory referenced by an address in a register
- ▶ g[X] for shared memory
- ▶ c[X][Y] for constant memory
- ▶ local[X] for local memory

Table 4 lists valid instructions for the GT200 GPUs.

Table 4 GT200 Instruction Set

Opcode	Description
A2R	Move address register to data register
ADA	Add immediate to address register
BAR	CTA-wide barrier synchronization
BRA	Conditional branch
BRK	Conditional break from a loop
BRX	Fetch an address from constant memory and branch to it

Opcode	Description
C2R	Conditional code to data register
CAL	Unconditional subroutine call
COS	Cosine
DADD	Double-precision floating point addition
DFMA	Double-precision floating point fused multiply-add
DMAX	Double-precision floating point maximum
DMIN	Double-precision floating point minimum
DMUL	Double-precision floating point multiply
DSET	Double-precision floating point conditional set
EX2	Exponential base two function
F2F	Copy floating-point value with conversion to a different floating-point type
F2I	Copy floating-point value with conversion to integer
FADD/FADD32/FADD32I	Single-precision floating point addition
FCMP	Single-precision floating point compare
FMAD/FMAD32/FMAD32I	Single-precision floating point multiply-add
FMAX	Single-precision floating point maximum
FMIN	Single-precision floating point minimum
FMUL/FMUL32/FMUL32I	Single-precision floating point multiply
FSET	Single-precision floating point conditional set
G2R	Move from shared memory to register. A .LCK suffix indicates that the bank is locked until a R2G.UNL has been performed; this is used to implement shared memory atomics.
GATOM.IADD/EXCH/CAS/IMIN/IMAX/INC/DEC/IAND/IOR/IXOR	Global memory atomic operations; performs both an atomic operation and returns the original value
GLD	Load from global memory
GRED.IADD/IMIN/IMAX/INC/DEC/IAND/IOR/IXOR	Global memory reduction operations; performs only an atomic operation with no return value
GST	Store to global memory
I2F	Copy integer value to floating-point with conversion
I2I	Copy integer value to integer with conversion
IADD/IADD32/IADD32I	Integer addition
IMAD/IMAD32/IMAD32I	Integer multiply-add
IMAX	Integer maximum
IMIN	Integer minimum
IMUL/IMUL32/IMUL32I	Integer multiply
ISAD/ISAD32	Sum of absolute difference

Opcode	Description
ISSET	Integer conditional set
LG2	Floating point logarithm base 2
LLD	Load from local memory
LST	Store to local memory
LOP	Logical operation (AND/OR/XOR)
MOV/MOV32	Move source to destination
MVC	Move from constant memory to destination
MVI	Move immediate to destination
NOP	No operation
R2A	Move register to address register
R2C	Move data register to conditional code
R2G	Store to shared memory. When used with the .UNL suffix, releases a previously held lock on that shared memory bank
RCP	Single-precision floating point reciprocal
RET	Conditional return from subroutine
RRO	Range reduction operator
RSQ	Reciprocal square root
S2R	Move special register to register
SHL	Shift left
SHR	Shift right
SIN	Sine
SSY	Set synchronization point; used before potentially divergent instructions
TEX/TEX32	Texture fetch
VOTE	Warp-vote primitive

4.2 Fermi Instruction Set

The Fermi architecture (Compute Capability 2.x) has the following instruction set format:

```
(instruction) (destination) (source1), (source2) ...
```

Valid destination and source locations include:

- ▶ RX for registers
- ▶ SRX for special system-controlled register
- ▶ PX for condition register
- ▶ c[X][Y] for constant memory

Table 5 lists valid instructions for the Fermi GPUs.

Table 5 Fermi Instruction Set

Opcode	Description
Floating Point Instructions	
FFMA	FP32 Fused Multiply Add
FADD	FP32 Add
FCMP	FP32 Compare
FMUL	FP32 Multiply
FMNMX	FP32 Minimum/Maximum
FSWZ	FP32 Swizzle
FSET	FP32 Set
FSETP	FP32 Set Predicate
RRO	FP Range Reduction Operator
MUFU	FP Multi-Function Operator
DFMA	FP64 Fused Multiply Add
DADD	FP64 Add
DMUL	FP64 Multiply
DMNMX	FP64 Minimum/Maximum
DSET	FP64 Set
DSETP	FP64 Set Predicate
Integer Instructions	
IMAD	Integer Multiply Add
IMUL	Integer Multiply
IADD	Integer Add
ISCADD	Integer Scaled Add
ISAD	Integer Sum Of Abs Diff
IMNMX	Integer Minimum/Maximum
BFE	Integer Bit Field Extract
BFI	Integer Bit Field Insert
SHR	Integer Shift Right
SHL	Integer Shift Left
LOP	Integer Logic Op
FLO	Integer Find Leading One
ISSET	Integer Set
ISSETP	Integer Set Predicate

Opcode	Description
ICMP	Integer Compare and Select
POPC	Population Count
Conversion Instructions	
F2F	Float to Float
F2I	Float to Integer
I2F	Integer to Float
I2I	Integer to Integer
Movement Instructions	
MOV	Move
SEL	Conditional Select/Move
PRMT	Permute
Predicate/CC Instructions	
P2R	Predicate to Register
R2P	Register to Predicate
CSET	CC Set
CSETP	CC Set Predicate
PSET	Predicate Set
PSETP	Predicate Set Predicate
Texture Instructions	
TEX	Texture Fetch
TLD	Texture Load
TLD4	Texture Load 4 Texels
TXQ	Texture Query
Compute Load/Store Instructions	
LDC	Load from Constant
LD	Load from Memory
LDU	Load Uniform
LDL	Load from Local Memory
LDS	Load from Shared Memory
LDLK	Load and Lock
LDSLK	Load from Shared Memory and Lock
LD_LDU	LD_LDU is a combination of a generic load LD with a load uniform LDU
LDS_LDU	LDS_LDU is combination of a Shared window load LDS with a load uniform LDU.
ST	Store to Memory

Opcode	Description
STL	Store to Local Memory
STUL	Store and Unlock
STS	Store to Shared Memory
STSUL	Store to Shared Memory and Unlock
ATOM	Atomic Memory Operation
RED	Atomic Memory Reduction Operation
CCTL	Cache Control
CCTL	Cache Control (Local)
MEMBAR	Memory Barrier
Surface	Memory Instructions
SULD	Surface Load
SULEA	Surface Load Effective Address
SUST	Surface Store
SURED	Surface Reduction
SUQ	Surface Query
Control Instructions	
BRA	Branch to Relative Address
BRX	Branch to Relative Indexed Address
JMP	Jump to Absolute Address
JMX	Jump to Absolute Indexed Address
CAL	Call to Relative Address
JCAL	Call to Absolute Address
RET	Return from Call
BRK	Break from Loop
CONT	Continue in Loop
LONGJMP	Long Jump
SSY	Set Sync Relative Address
PBK	Pre-Break Relative Address
PCNT	Pre-Continue Relative Address
PRET	Pre-Return Relative Address
PLONGJMP	Pre-Long-Jump Relative Address
BPT	Breakpoint/Trap
EXIT	Exit Program
Miscellaneous Instructions	
NOP	No Operation

Opcode	Description
S2R	Special Register to Register
B2R	Barrier to Register
LEPC	Load Effective PC
BAR	Barrier Synchronization
VOTE	Query condition across threads

4.3 Kepler Instruction Set

The Kepler architecture (Compute Capability 3.x) has the following instruction set format:

```
(instruction) (destination) (source1), (source2) ...
```

Valid destination and source locations include:

- ▶ X for registers
- ▶ SRX for special system-controlled registers
- ▶ PX for condition registers
- ▶ c[X][Y] for constant memory

Table 6 lists valid instructions for the Kepler GPUs.

Table 6 Kepler Instruction Set

Opcode	Description
Floating Point Instructions	
FFMA	FP32 Fused Multiply Add
FADD	FP32 Add
FCMP	FP32 Compare
FMUL	FP32 Multiply
FMNMX	FP32 Minimum/Maximum
FSWZ	FP32 Swizzle
FSET	FP32 Set
FSETP	FP32 Set Predicate
FCHK	FP32 Division Test
RRO	FP Range Reduction Operator
MUFU	FP Multi-Function Operator
DFMA	FP64 Fused Multiply Add
DADD	FP64 Add
DMUL	FP64 Multiply

Opcode	Description
DMNMX	FP64 Minimum/Maximum
DSET	FP64 Set
DSETP	FP64 Set Predicate
Integer Instructions	
IMAD	Integer Multiply Add
IMADSP	Integer Extract Multiply Add
IMUL	Integer Multiply
IADD	Integer Add
ISCADD	Integer Scaled Add
ISAD	Integer Sum Of Abs Diff
IMNMX	Integer Minimum/Maximum
BFE	Integer Bit Field Extract
BFI	Integer Bit Field Insert
SHR	Integer Shift Right
SHL	Integer Shift Left
SHF	Integer Funnel Shift
LOP	Integer Logic Op
FLO	Integer Find Leading One
ISSET	Integer Set
ISETP	Integer Set Predicate
ICMP	Integer Compare and Select
POPC	Population Count
Conversion Instructions	
F2F	Float to Float
F2I	Float to Integer
I2F	Integer to Float
I2I	Integer to Integer
Movement Instructions	
MOV	Move
SEL	Conditional Select/Move
PRMT	Permute
SHFL	Warp Shuffle
Predicate/CC Instructions	
P2R	Predicate to Register
R2P	Register to Predicate

Opcode	Description
CSET	CC Set
CSETP	CC Set Predicate
PSET	Predicate Set
PSETP	Predicate Set Predicate
Texture Instructions	
TEX	Texture Fetch
TLD	Texture Load
TLD4	Texture Load 4 Texels
TXQ	Texture Query
Compute Load/Store Instructions	
LDC	Load from Constant
LD	Load from Memory
LDG	Non-coherent Global Memory Load
LDL	Load from Local Memory
LDS	Load from Shared Memory
LDSLK	Load from Shared Memory and Lock
ST	Store to Memory
STL	Store to Local Memory
STS	Store to Shared Memory
STSCUL	Store to Shared Memory Conditionally and Unlock
ATOM	Atomic Memory Operation
RED	Atomic Memory Reduction Operation
CCTL	Cache Control
CCTLL	Cache Control (Local)
MEMBAR	Memory Barrier
Surface	Memory Instructions
SUCLAMP	Surface Clamp
SUBFM	Surface Bit Field Merge
SUEAU	Surface Effective Address
SULDGA	Surface Load Generic Address
SUSTGA	Surface Store Generic Address
Control Instructions	
BRA	Branch to Relative Address
BRX	Branch to Relative Indexed Address
JMP	Jump to Absolute Address

Opcode	Description
JMX	Jump to Absolute Indexed Address
CAL	Call to Relative Address
JCAL	Call to Absolute Address
RET	Return from Call
BRK	Break from Loop
CONT	Continue in Loop
SSY	Set Sync Relative Address
PBK	Pre-Break Relative Address
PCNT	Pre-Continue Relative Address
PRET	Pre-Return Relative Address
BPT	Breakpoint/Trap
EXIT	Exit Program
Miscellaneous Instructions	
NOP	No Operation
S2R	Special Register to Register
B2R	Barrier to Register
BAR	Barrier Synchronization
VOTE	Query condition across threads

Notice

ALL NVIDIA DESIGN SPECIFICATIONS, REFERENCE BOARDS, FILES, DRAWINGS, DIAGNOSTICS, LISTS, AND OTHER DOCUMENTS (TOGETHER AND SEPARATELY, "MATERIALS") ARE BEING PROVIDED "AS IS." NVIDIA MAKES NO WARRANTIES, EXPRESSED, IMPLIED, STATUTORY, OR OTHERWISE WITH RESPECT TO THE MATERIALS, AND EXPRESSLY DISCLAIMS ALL IMPLIED WARRANTIES OF NONINFRINGEMENT, MERCHANTABILITY, AND FITNESS FOR A PARTICULAR PURPOSE.

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

Trademarks

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

Copyright

© 2013 NVIDIA Corporation. All rights reserved.