

**NAME**

Data types used by CUDA driver –

**Data Structures**

```

struct CUDA_ARRAY3D_DESCRIPTOR
struct CUDA_ARRAY_DESCRIPTOR
struct CUDA_MEMCPY2D
struct CUDA_MEMCPY3D
struct CUDA_MEMCPY3D_PEER
struct CUDA_POINTER_ATTRIBUTE_P2P_TOKENS
struct CUDA_RESOURCE_DESC
struct CUDA_RESOURCE_VIEW_DESC
struct CUDA_TEXTURE_DESC
struct CUdevprop
struct CUipcEventHandle
struct CUipcMemHandle

```

**Defines**

```

#define CU_IPC_HANDLE_SIZE 64
#define CU_LAUNCH_PARAM_BUFFER_POINTER ((void*)0x01)
#define CU_LAUNCH_PARAM_BUFFER_SIZE ((void*)0x02)
#define CU_LAUNCH_PARAM_END ((void*)0x00)
#define CU_MEMHOSTALLOC_DEVICEMAP 0x02
#define CU_MEMHOSTALLOC_PORTABLE 0x01
#define CU_MEMHOSTALLOC_WRITECOMBINED 0x04
#define CU_MEMHOSTREGISTER_DEVICEMAP 0x02
#define CU_MEMHOSTREGISTER_PORTABLE 0x01
#define CU_PARAM_TR_DEFAULT -1
#define CU_TRSA_OVERRIDE_FORMAT 0x01
#define CU_TRSF_NORMALIZED_COORDINATES 0x02
#define CU_TRSF_READ_AS_INTEGER 0x01
#define CU_TRSF_SRGB 0x10
#define CUDA_ARRAY3D_2DARRAY 0x01
#define CUDA_ARRAY3D_CUBEMAP 0x04
#define CUDA_ARRAY3D_DEPTH_TEXTURE 0x10
#define CUDA_ARRAY3D_LAYERED 0x01
#define CUDA_ARRAY3D_SURFACE_LDST 0x02
#define CUDA_ARRAY3D_TEXTURE_GATHER 0x08
#define CUDA_VERSION 6000

```

**Typedefs**

```

typedef struct CUarray_st * CUarray
typedef struct CUctx_st * CUcontext
typedef int CUdevice
typedef unsigned int CUdeviceptr
typedef struct CUevent_st * CUevent
typedef struct CUfunc_st * CUfunction
typedef struct CUgraphicsResource_st * CUgraphicsResource
typedef struct CUmipmappedArray_st * CUmipmappedArray
typedef struct CUmod_st * CUmodule
typedef struct CUstream_st * CUstream
typedef void(CUDA_CB * CUstreamCallback)(CUstream hStream, CUresult status, void
    *userData)
typedef unsigned long long CUsurfObject
typedef struct CUsurfref_st * CUsurfref
typedef unsigned long long CUtexObject
typedef struct CUtexref_st * CUtexref

```

**Enumerations**

```

enum CUaddress_mode { CU_TR_ADDRESS_MODE_WRAP = 0,
    CU_TR_ADDRESS_MODE_CLAMP = 1, CU_TR_ADDRESS_MODE_MIRROR = 2,

```



```

    CU_TR_ADDRESS_MODE_BORDER = 3 }
enum CUarray_cubemap_face { CU_CUBEMAP_FACE_POSITIVE_X = 0x00,
    CU_CUBEMAP_FACE_NEGATIVE_X = 0x01, CU_CUBEMAP_FACE_POSITIVE_Y =
    0x02, CU_CUBEMAP_FACE_NEGATIVE_Y = 0x03,
    CU_CUBEMAP_FACE_POSITIVE_Z = 0x04, CU_CUBEMAP_FACE_NEGATIVE_Z =
    0x05 }
enum CUarray_format { CU_AD_FORMAT_UNSIGNED_INT8 = 0x01,
    CU_AD_FORMAT_UNSIGNED_INT16 = 0x02, CU_AD_FORMAT_UNSIGNED_INT32
    = 0x03, CU_AD_FORMAT_SIGNED_INT8 = 0x08, CU_AD_FORMAT_SIGNED_INT16
    = 0x09, CU_AD_FORMAT_SIGNED_INT32 = 0x0a, CU_AD_FORMAT_HALF = 0x10,
    CU_AD_FORMAT_FLOAT = 0x20 }
enum CUcomputemode { CU_COMPUTEMODE_DEFAULT = 0,
    CU_COMPUTEMODE_EXCLUSIVE = 1, CU_COMPUTEMODE_PROHIBITED = 2,
    CU_COMPUTEMODE_EXCLUSIVE_PROCESS = 3 }
enum CUctx_flags { CU_CTX_SCHED_AUTO = 0x00, CU_CTX_SCHED_SPIN = 0x01,
    CU_CTX_SCHED_YIELD = 0x02, CU_CTX_SCHED_BLOCKING_SYNC = 0x04,
    CU_CTX_BLOCKING_SYNC = 0x04, CU_CTX_MAP_HOST = 0x08,
    CU_CTX_LMEM_RESIZE_TO_MAX = 0x10 }
enum CUdevice_attribute { CU_DEVICE_ATTRIBUTE_MAX_THREADS_PER_BLOCK = 1,
    CU_DEVICE_ATTRIBUTE_MAX_BLOCK_DIM_X = 2,
    CU_DEVICE_ATTRIBUTE_MAX_BLOCK_DIM_Y = 3,
    CU_DEVICE_ATTRIBUTE_MAX_BLOCK_DIM_Z = 4,
    CU_DEVICE_ATTRIBUTE_MAX_GRID_DIM_X = 5,
    CU_DEVICE_ATTRIBUTE_MAX_GRID_DIM_Y = 6,
    CU_DEVICE_ATTRIBUTE_MAX_GRID_DIM_Z = 7,
    CU_DEVICE_ATTRIBUTE_MAX_SHARED_MEMORY_PER_BLOCK = 8,
    CU_DEVICE_ATTRIBUTE_SHARED_MEMORY_PER_BLOCK = 8,
    CU_DEVICE_ATTRIBUTE_TOTAL_CONSTANT_MEMORY = 9,
    CU_DEVICE_ATTRIBUTE_WARP_SIZE = 10,
    CU_DEVICE_ATTRIBUTE_MAX_PITCH = 11,
    CU_DEVICE_ATTRIBUTE_MAX_REGISTERS_PER_BLOCK = 12,
    CU_DEVICE_ATTRIBUTE_REGISTERS_PER_BLOCK = 12,
    CU_DEVICE_ATTRIBUTE_CLOCK_RATE = 13,
    CU_DEVICE_ATTRIBUTE_TEXTURE_ALIGNMENT = 14,
    CU_DEVICE_ATTRIBUTE_GPU_OVERLAP = 15,
    CU_DEVICE_ATTRIBUTE_MULTIPROCESSOR_COUNT = 16,
    CU_DEVICE_ATTRIBUTE_KERNEL_EXEC_TIMEOUT = 17,
    CU_DEVICE_ATTRIBUTE_INTEGRATED = 18,
    CU_DEVICE_ATTRIBUTE_CAN_MAP_HOST_MEMORY = 19,
    CU_DEVICE_ATTRIBUTE_COMPUTE_MODE = 20,
    CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE1D_WIDTH = 21,
    CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE2D_WIDTH = 22,
    CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE2D_HEIGHT = 23,
    CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE3D_WIDTH = 24,
    CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE3D_HEIGHT = 25,
    CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE3D_DEPTH = 26,
    CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE2D_LAYERED_WIDTH = 27,
    CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE2D_LAYERED_HEIGHT = 28,
    CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE2D_LAYERED_LAYERS = 29,
    CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE2D_ARRAY_WIDTH = 27,
    CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE2D_ARRAY_HEIGHT = 28,
    CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE2D_ARRAY_NUMSLICES = 29,
    CU_DEVICE_ATTRIBUTE_SURFACE_ALIGNMENT = 30,
    CU_DEVICE_ATTRIBUTE_CONCURRENT_KERNELS = 31,
    CU_DEVICE_ATTRIBUTE_ECC_ENABLED = 32,
    CU_DEVICE_ATTRIBUTE_PCI_BUS_ID = 33,
    CU_DEVICE_ATTRIBUTE_PCI_DEVICE_ID = 34,
    CU_DEVICE_ATTRIBUTE_TCC_DRIVER = 35,
    CU_DEVICE_ATTRIBUTE_MEMORY_CLOCK_RATE = 36,

```



```

CU_DEVICE_ATTRIBUTE_GLOBAL_MEMORY_BUS_WIDTH = 37,
CU_DEVICE_ATTRIBUTE_L2_CACHE_SIZE = 38,
CU_DEVICE_ATTRIBUTE_MAX_THREADS_PER_MULTIPROCESSOR = 39,
CU_DEVICE_ATTRIBUTE_ASYNC_ENGINE_COUNT = 40,
CU_DEVICE_ATTRIBUTE_UNIFIED_ADDRESSING = 41,
CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE1D_LAYERED_WIDTH = 42,
CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE1D_LAYERED_LAYERS = 43,
CU_DEVICE_ATTRIBUTE_CAN_TEX2D_GATHER = 44,
CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE2D_GATHER_WIDTH = 45,
CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE2D_GATHER_HEIGHT = 46,
CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE3D_WIDTH_ALTERNATE = 47,
CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE3D_HEIGHT_ALTERNATE = 48,
CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE3D_DEPTH_ALTERNATE = 49,
CU_DEVICE_ATTRIBUTE_PCI_DOMAIN_ID = 50,
CU_DEVICE_ATTRIBUTE_TEXTURE_PITCH_ALIGNMENT = 51,
CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURECUBEMAP_WIDTH = 52,
CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURECUBEMAP_LAYERED_WIDTH =
53,
CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURECUBEMAP_LAYERED_LAYERS
= 54, CU_DEVICE_ATTRIBUTE_MAXIMUM_SURFACE1D_WIDTH = 55,
CU_DEVICE_ATTRIBUTE_MAXIMUM_SURFACE2D_WIDTH = 56,
CU_DEVICE_ATTRIBUTE_MAXIMUM_SURFACE2D_HEIGHT = 57,
CU_DEVICE_ATTRIBUTE_MAXIMUM_SURFACE3D_WIDTH = 58,
CU_DEVICE_ATTRIBUTE_MAXIMUM_SURFACE3D_HEIGHT = 59,
CU_DEVICE_ATTRIBUTE_MAXIMUM_SURFACE3D_DEPTH = 60,
CU_DEVICE_ATTRIBUTE_MAXIMUM_SURFACE1D_LAYERED_WIDTH = 61,
CU_DEVICE_ATTRIBUTE_MAXIMUM_SURFACE1D_LAYERED_LAYERS = 62,
CU_DEVICE_ATTRIBUTE_MAXIMUM_SURFACE2D_LAYERED_WIDTH = 63,
CU_DEVICE_ATTRIBUTE_MAXIMUM_SURFACE2D_LAYERED_HEIGHT = 64,
CU_DEVICE_ATTRIBUTE_MAXIMUM_SURFACE2D_LAYERED_LAYERS = 65,
CU_DEVICE_ATTRIBUTE_MAXIMUM_SURFACECUBEMAP_WIDTH = 66,
CU_DEVICE_ATTRIBUTE_MAXIMUM_SURFACECUBEMAP_LAYERED_WIDTH =
67,
CU_DEVICE_ATTRIBUTE_MAXIMUM_SURFACECUBEMAP_LAYERED_LAYERS
= 68, CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE1D_LINEAR_WIDTH = 69,
CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE2D_LINEAR_WIDTH = 70,
CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE2D_LINEAR_HEIGHT = 71,
CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE2D_LINEAR_PITCH = 72,
CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE2D_MIPMAPPED_WIDTH = 73,
CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE2D_MIPMAPPED_HEIGHT = 74,
CU_DEVICE_ATTRIBUTE_COMPUTE_CAPABILITY_MAJOR = 75,
CU_DEVICE_ATTRIBUTE_COMPUTE_CAPABILITY_MINOR = 76,
CU_DEVICE_ATTRIBUTE_MAXIMUM_TEXTURE1D_MIPMAPPED_WIDTH = 77,
CU_DEVICE_ATTRIBUTE_STREAM_PRIORITIES_SUPPORTED = 78,
CU_DEVICE_ATTRIBUTE_GLOBAL_L1_CACHE_SUPPORTED = 79,
CU_DEVICE_ATTRIBUTE_LOCAL_L1_CACHE_SUPPORTED = 80,
CU_DEVICE_ATTRIBUTE_MAX_SHARED_MEMORY_PER_MULTIPROCESSOR =
81, CU_DEVICE_ATTRIBUTE_MAX_REGISTERS_PER_MULTIPROCESSOR = 82,
CU_DEVICE_ATTRIBUTE_MANAGED_MEMORY = 83,
CU_DEVICE_ATTRIBUTE_MULTI_GPU_BOARD = 84,
CU_DEVICE_ATTRIBUTE_MULTI_GPU_BOARD_GROUP_ID = 85 }
enum CUevent_flags { CU_EVENT_DEFAULT = 0x0, CU_EVENT_BLOCKING_SYNC = 0x1,
CU_EVENT_DISABLE_TIMING = 0x2, CU_EVENT_INTERPROCESS = 0x4 }
enum CUfilter_mode { CU_TR_FILTER_MODE_POINT = 0,
CU_TR_FILTER_MODE_LINEAR = 1 }
enum CUfunc_cache { CU_FUNC_CACHE_PREFER_NONE = 0x00,
CU_FUNC_CACHE_PREFER_SHARED = 0x01, CU_FUNC_CACHE_PREFER_L1 =
0x02, CU_FUNC_CACHE_PREFER_EQUAL = 0x03 }
enum CUfunction_attribute { CU_FUNC_ATTRIBUTE_MAX_THREADS_PER_BLOCK = 0,

```



```

    CU_FUNC_ATTRIBUTE_SHARED_SIZE_BYTES = 1,
    CU_FUNC_ATTRIBUTE_CONST_SIZE_BYTES = 2,
    CU_FUNC_ATTRIBUTE_LOCAL_SIZE_BYTES = 3,
    CU_FUNC_ATTRIBUTE_NUM_REGS = 4, CU_FUNC_ATTRIBUTE_PTX_VERSION
    = 5, CU_FUNC_ATTRIBUTE_BINARY_VERSION = 6,
    CU_FUNC_ATTRIBUTE_CACHE_MODE_CA = 7 }
enum CUgraphicsMapResourceFlags
enum CUgraphicsRegisterFlags
enum CUipcMem_flags { CU_IPC_MEM_LAZY_ENABLE_PEER_ACCESS = 0x1 }
enum CUjit_cacheMode { CU_JIT_CACHE_OPTION_NONE = 0,
    CU_JIT_CACHE_OPTION_CG, CU_JIT_CACHE_OPTION_CA }
enum CUjit_fallback { CU_PREFER_PTX = 0, CU_PREFER_BINARY }
enum CUjit_option { CU_JIT_MAX_REGISTERS = 0, CU_JIT_THREADS_PER_BLOCK,
    CU_JIT_WALL_TIME, CU_JIT_INFO_LOG_BUFFER,
    CU_JIT_INFO_LOG_BUFFER_SIZE_BYTES, CU_JIT_ERROR_LOG_BUFFER,
    CU_JIT_ERROR_LOG_BUFFER_SIZE_BYTES, CU_JIT_OPTIMIZATION_LEVEL,
    CU_JIT_TARGET_FROM_CUCONTEXT, CU_JIT_TARGET,
    CU_JIT_FALLBACK_STRATEGY, CU_JIT_GENERATE_DEBUG_INFO,
    CU_JIT_LOG_VERBOSE, CU_JIT_GENERATE_LINE_INFO,
    CU_JIT_CACHE_MODE }
enum CUjit_target { CU_TARGET_COMPUTE_10 = 10, CU_TARGET_COMPUTE_11 = 11,
    CU_TARGET_COMPUTE_12 = 12, CU_TARGET_COMPUTE_13 = 13,
    CU_TARGET_COMPUTE_20 = 20, CU_TARGET_COMPUTE_21 = 21,
    CU_TARGET_COMPUTE_30 = 30, CU_TARGET_COMPUTE_32 = 32,
    CU_TARGET_COMPUTE_35 = 35, CU_TARGET_COMPUTE_50 = 50 }
enum CUjitInputType { CU_JIT_INPUT_CUBIN = 0, CU_JIT_INPUT_PTX,
    CU_JIT_INPUT_FATBINARY, CU_JIT_INPUT_OBJECT, CU_JIT_INPUT_LIBRARY
    }
enum CULimit { CU_LIMIT_STACK_SIZE = 0x00, CU_LIMIT_PRINTF_FIFO_SIZE = 0x01,
    CU_LIMIT_MALLOC_HEAP_SIZE = 0x02,
    CU_LIMIT_DEV_RUNTIME_SYNC_DEPTH = 0x03,
    CU_LIMIT_DEV_RUNTIME_PENDING_LAUNCH_COUNT = 0x04 }
enum CUmемAttach_flags { CU_MEM_ATTACH_GLOBAL = 0x1,
    CU_MEM_ATTACH_HOST = 0x2, CU_MEM_ATTACH_SINGLE = 0x4 }
enum CUmemorytype { CU_MEMORYTYPE_HOST = 0x01, CU_MEMORYTYPE_DEVICE =
    0x02, CU_MEMORYTYPE_ARRAY = 0x03, CU_MEMORYTYPE_UNIFIED = 0x04 }
enum CUpointer_attribute { CU_POINTER_ATTRIBUTE_CONTEXT = 1,
    CU_POINTER_ATTRIBUTE_MEMORY_TYPE = 2,
    CU_POINTER_ATTRIBUTE_DEVICE_POINTER = 3,
    CU_POINTER_ATTRIBUTE_HOST_POINTER = 4,
    CU_POINTER_ATTRIBUTE_P2P_TOKENS = 5,
    CU_POINTER_ATTRIBUTE_SYNC_MEMOPS = 6,
    CU_POINTER_ATTRIBUTE_BUFFER_ID = 7,
    CU_POINTER_ATTRIBUTE_IS_MANAGED = 8 }
enum CUresourcetype { CU_RESOURCE_TYPE_ARRAY = 0x00,
    CU_RESOURCE_TYPE_MIPMAPPED_ARRAY = 0x01,
    CU_RESOURCE_TYPE_LINEAR = 0x02, CU_RESOURCE_TYPE_PITCH2D = 0x03 }
enum CUresourceViewFormat { CU_RES_VIEW_FORMAT_NONE = 0x00,
    CU_RES_VIEW_FORMAT_UINT_1X8 = 0x01, CU_RES_VIEW_FORMAT_UINT_2X8
    = 0x02, CU_RES_VIEW_FORMAT_UINT_4X8 = 0x03,
    CU_RES_VIEW_FORMAT_SINT_1X8 = 0x04, CU_RES_VIEW_FORMAT_SINT_2X8
    = 0x05, CU_RES_VIEW_FORMAT_SINT_4X8 = 0x06,
    CU_RES_VIEW_FORMAT_UINT_1X16 = 0x07,
    CU_RES_VIEW_FORMAT_UINT_2X16 = 0x08,
    CU_RES_VIEW_FORMAT_UINT_4X16 = 0x09,
    CU_RES_VIEW_FORMAT_SINT_1X16 = 0x0a,
    CU_RES_VIEW_FORMAT_SINT_2X16 = 0x0b,
    CU_RES_VIEW_FORMAT_SINT_4X16 = 0x0c,
    CU_RES_VIEW_FORMAT_UINT_1X32 = 0x0d,

```



```

CU_RES_VIEW_FORMAT_UINT_2X32 = 0x0e,
CU_RES_VIEW_FORMAT_UINT_4X32 = 0x0f,
CU_RES_VIEW_FORMAT_SINT_1X32 = 0x10,
CU_RES_VIEW_FORMAT_SINT_2X32 = 0x11,
CU_RES_VIEW_FORMAT_SINT_4X32 = 0x12,
CU_RES_VIEW_FORMAT_FLOAT_1X16 = 0x13,
CU_RES_VIEW_FORMAT_FLOAT_2X16 = 0x14,
CU_RES_VIEW_FORMAT_FLOAT_4X16 = 0x15,
CU_RES_VIEW_FORMAT_FLOAT_1X32 = 0x16,
CU_RES_VIEW_FORMAT_FLOAT_2X32 = 0x17,
CU_RES_VIEW_FORMAT_FLOAT_4X32 = 0x18,
CU_RES_VIEW_FORMAT_UNSIGNED_BC1 = 0x19,
CU_RES_VIEW_FORMAT_UNSIGNED_BC2 = 0x1a,
CU_RES_VIEW_FORMAT_UNSIGNED_BC3 = 0x1b,
CU_RES_VIEW_FORMAT_UNSIGNED_BC4 = 0x1c,
CU_RES_VIEW_FORMAT_SIGNED_BC4 = 0x1d,
CU_RES_VIEW_FORMAT_UNSIGNED_BC5 = 0x1e,
CU_RES_VIEW_FORMAT_SIGNED_BC5 = 0x1f,
CU_RES_VIEW_FORMAT_UNSIGNED_BC6H = 0x20,
CU_RES_VIEW_FORMAT_SIGNED_BC6H = 0x21,
CU_RES_VIEW_FORMAT_UNSIGNED_BC7 = 0x22 }
enum CudaResult { CUDA_SUCCESS = 0, CUDA_ERROR_INVALID_VALUE = 1,
    CUDA_ERROR_OUT_OF_MEMORY = 2, CUDA_ERROR_NOT_INITIALIZED = 3,
    CUDA_ERROR_DEINITIALIZED = 4, CUDA_ERROR_PROFILER_DISABLED = 5,
    CUDA_ERROR_PROFILER_NOT_INITIALIZED = 6,
    CUDA_ERROR_PROFILER_ALREADY_STARTED = 7,
    CUDA_ERROR_PROFILER_ALREADY_STOPPED = 8,
    CUDA_ERROR_NO_DEVICE = 100, CUDA_ERROR_INVALID_DEVICE = 101,
    CUDA_ERROR_INVALID_IMAGE = 200, CUDA_ERROR_INVALID_CONTEXT =
    201, CUDA_ERROR_CONTEXT_ALREADY_CURRENT = 202,
    CUDA_ERROR_MAP_FAILED = 205, CUDA_ERROR_UNMAP_FAILED = 206,
    CUDA_ERROR_ARRAY_IS_MAPPED = 207, CUDA_ERROR_ALREADY_MAPPED =
    208, CUDA_ERROR_NO_BINARY_FOR_GPU = 209,
    CUDA_ERROR_ALREADY_ACQUIRED = 210, CUDA_ERROR_NOT_MAPPED =
    211, CUDA_ERROR_NOT_MAPPED_AS_ARRAY = 212,
    CUDA_ERROR_NOT_MAPPED_AS_POINTER = 213,
    CUDA_ERROR_ECC_UNCORRECTABLE = 214,
    CUDA_ERROR_UNSUPPORTED_LIMIT = 215,
    CUDA_ERROR_CONTEXT_ALREADY_IN_USE = 216,
    CUDA_ERROR_PEER_ACCESS_UNSUPPORTED = 217,
    CUDA_ERROR_INVALID_PTX = 218, CUDA_ERROR_INVALID_SOURCE = 300,
    CUDA_ERROR_FILE_NOT_FOUND = 301,
    CUDA_ERROR_SHARED_OBJECT_SYMBOL_NOT_FOUND = 302,
    CUDA_ERROR_SHARED_OBJECT_INIT_FAILED = 303,
    CUDA_ERROR_OPERATING_SYSTEM = 304, CUDA_ERROR_INVALID_HANDLE =
    400, CUDA_ERROR_NOT_FOUND = 500, CUDA_ERROR_NOT_READY = 600,
    CUDA_ERROR_ILLEGAL_ADDRESS = 700,
    CUDA_ERROR_LAUNCH_OUT_OF_RESOURCES = 701,
    CUDA_ERROR_LAUNCH_TIMEOUT = 702,
    CUDA_ERROR_LAUNCH_INCOMPATIBLE_TEXTURING = 703,
    CUDA_ERROR_PEER_ACCESS_ALREADY_ENABLED = 704,
    CUDA_ERROR_PEER_ACCESS_NOT_ENABLED = 705,
    CUDA_ERROR_PRIMARY_CONTEXT_ACTIVE = 708,
    CUDA_ERROR_CONTEXT_IS_DESTROYED = 709, CUDA_ERROR_ASSERT = 710,
    CUDA_ERROR_TOO_MANY_PEERS = 711,
    CUDA_ERROR_HOST_MEMORY_ALREADY_REGISTERED = 712,
    CUDA_ERROR_HOST_MEMORY_NOT_REGISTERED = 713,
    CUDA_ERROR_HARDWARE_STACK_ERROR = 714,
    CUDA_ERROR_ILLEGAL_INSTRUCTION = 715,

```



```

    CUDA_ERROR_MISALIGNED_ADDRESS = 716,
    CUDA_ERROR_INVALID_ADDRESS_SPACE = 717, CUDA_ERROR_INVALID_PC =
    718, CUDA_ERROR_LAUNCH_FAILED = 719, CUDA_ERROR_NOT_PERMITTED =
    800, CUDA_ERROR_NOT_SUPPORTED = 801, CUDA_ERROR_UNKNOWN = 999 }
enum CUsharedconfig { CU_SHARED_MEM_CONFIG_DEFAULT_BANK_SIZE = 0x00,
    CU_SHARED_MEM_CONFIG_FOUR_BYTE_BANK_SIZE = 0x01,
    CU_SHARED_MEM_CONFIG_EIGHT_BYTE_BANK_SIZE = 0x02 }
enum CUstream_flags { CU_STREAM_DEFAULT = 0x0, CU_STREAM_NON_BLOCKING =
    0x1 }

```

## Define Documentation

**#define CU\_IPC\_HANDLE\_SIZE 64**

CUDA IPC handle size

**#define CU\_LAUNCH\_PARAM\_BUFFER\_POINTER ((void\*)0x01)**

Indicator that the next value in the `extra` parameter to `cuLaunchKernel` will be a pointer to a buffer containing all kernel parameters used for launching kernel `f`. This buffer needs to honor all alignment/padding requirements of the individual parameters. If

`CU_LAUNCH_PARAM_BUFFER_SIZE` is not also specified in the `extra` array, then `CU_LAUNCH_PARAM_BUFFER_POINTER` will have no effect.

**#define CU\_LAUNCH\_PARAM\_BUFFER\_SIZE ((void\*)0x02)**

Indicator that the next value in the `extra` parameter to `cuLaunchKernel` will be a pointer to a `size_t` which contains the size of the buffer specified with `CU_LAUNCH_PARAM_BUFFER_POINTER`. It is required that `CU_LAUNCH_PARAM_BUFFER_POINTER` also be specified in the `extra` array if the value associated with `CU_LAUNCH_PARAM_BUFFER_SIZE` is not zero.

**#define CU\_LAUNCH\_PARAM\_END ((void\*)0x00)**

End of array terminator for the `extra` parameter to `cuLaunchKernel`

**#define CU\_MEMHOSTALLOC\_DEVICEMAP 0x02**

If set, host memory is mapped into CUDA address space and `cuMemHostGetDevicePointer()` may be called on the host pointer. Flag for `cuMemHostAlloc()`

**#define CU\_MEMHOSTALLOC\_PORTABLE 0x01**

If set, host memory is portable between CUDA contexts. Flag for `cuMemHostAlloc()`

**#define CU\_MEMHOSTALLOC\_WRITECOMBINED 0x04**

If set, host memory is allocated as write-combined - fast to write, faster to DMA, slow to read except via SSE4 streaming load instruction (MOVNTDQA). Flag for `cuMemHostAlloc()`

**#define CU\_MEMHOSTREGISTER\_DEVICEMAP 0x02**

If set, host memory is mapped into CUDA address space and `cuMemHostGetDevicePointer()` may be called on the host pointer. Flag for `cuMemHostRegister()`

**#define CU\_MEMHOSTREGISTER\_PORTABLE 0x01**

If set, host memory is portable between CUDA contexts. Flag for `cuMemHostRegister()`

**#define CU\_PARAM\_TR\_DEFAULT -1**

For texture references loaded into the module, use default texunit from texture reference.

**#define CU\_TRSA\_OVERRIDE\_FORMAT 0x01**

Override the texref format with a format inferred from the array. Flag for `cuTexRefSetArray()`

**#define CU\_TRSF\_NORMALIZED\_COORDINATES 0x02**

Use normalized texture coordinates in the range [0,1) instead of [0,dim). Flag for `cuTexRefSetFlags()`

**#define CU\_TRSF\_READ\_AS\_INTEGER 0x01**

Read the texture as integers rather than promoting the values to floats in the range [0,1]. Flag for `cuTexRefSetFlags()`

**#define CU\_TRSF\_SRGB 0x10**

Perform sRGB->linear conversion during texture read. Flag for `cuTexRefSetFlags()`

**#define CUDA\_ARRAY3D\_2DARRAY 0x01**

Deprecated, use `CUDA_ARRAY3D_LAYERED`



**#define CUDA\_ARRAY3D\_CUBEMAP 0x04**

If set, the CUDA array is a collection of six 2D arrays, representing faces of a cube. The width of such a CUDA array must be equal to its height, and Depth must be six. If **CUDA\_ARRAY3D\_LAYERED** flag is also set, then the CUDA array is a collection of cubemaps and Depth must be a multiple of six.

**#define CUDA\_ARRAY3D\_DEPTH\_TEXTURE 0x10**

This flag if set indicates that the CUDA array is a **DEPTH\_TEXTURE**.

**#define CUDA\_ARRAY3D\_LAYERED 0x01**

If set, the CUDA array is a collection of layers, where each layer is either a 1D or a 2D array and the Depth member of **CUDA\_ARRAY3D\_DESCRIPTOR** specifies the number of layers, not the depth of a 3D array.

**#define CUDA\_ARRAY3D\_SURFACE\_LDST 0x02**

This flag must be set in order to bind a surface reference to the CUDA array

**#define CUDA\_ARRAY3D\_TEXTURE\_GATHER 0x08**

This flag must be set in order to perform texture gather operations on a CUDA array.

**#define CUDA\_VERSION 6000**

CUDA API version number

**Typedef Documentation****typedef struct CUarray\_st\* CUarray**

CUDA array

**typedef struct CUctx\_st\* CUcontext**

CUDA context

**typedef int CUdevice**

CUDA device

**typedef unsigned int CUdeviceptr**

CUDA device pointer

**typedef struct CUevent\_st\* CUevent**

CUDA event

**typedef struct CUfunc\_st\* CUfunction**

CUDA function

**typedef struct CUgraphicsResource\_st\* CUgraphicsResource**

CUDA graphics interop resource

**typedef struct CUmipmappedArray\_st\* CUmipmappedArray**

CUDA mipmapped array

**typedef struct CUmod\_st\* CUmodule**

CUDA module

**typedef struct CUstream\_st\* CUstream**

CUDA stream

**typedef void(CUDA\_CB \* CUstreamCallback)(CUstream hStream, CUresult status, void \*userData)**

CUDA stream callback

**Parameters:**

*hStream* The stream the callback was added to, as passed to **cuStreamAddCallback**. May be NULL.

*status* **CUDA\_SUCCESS** or any persistent error on the stream.

*userData* User parameter provided at registration.

**typedef unsigned long long CUsurfObject**

CUDA surface object

**typedef struct CUsurfref\_st\* CUsurfref**

CUDA surface reference

**typedef unsigned long long CUTexObject**

CUDA texture object



**typedef struct CUtexref\_st\* CUtexref**

CUDA texture reference

## Enumeration Type Documentation

**enum CUaddress\_mode**

Texture reference addressing modes

**Enumerator:**

*CU\_TR\_ADDRESS\_MODE\_WRAP*

Wrapping address mode

*CU\_TR\_ADDRESS\_MODE\_CLAMP*

Clamp to edge address mode

*CU\_TR\_ADDRESS\_MODE\_MIRROR*

Mirror address mode

*CU\_TR\_ADDRESS\_MODE\_BORDER*

Border address mode

**enum CUarray\_cubemap\_face**

Array indices for cube faces

**Enumerator:**

*CU\_CUBEMAP\_FACE\_POSITIVE\_X*

Positive X face of cubemap

*CU\_CUBEMAP\_FACE\_NEGATIVE\_X*

Negative X face of cubemap

*CU\_CUBEMAP\_FACE\_POSITIVE\_Y*

Positive Y face of cubemap

*CU\_CUBEMAP\_FACE\_NEGATIVE\_Y*

Negative Y face of cubemap

*CU\_CUBEMAP\_FACE\_POSITIVE\_Z*

Positive Z face of cubemap

*CU\_CUBEMAP\_FACE\_NEGATIVE\_Z*

Negative Z face of cubemap

**enum CUarray\_format**

Array formats

**Enumerator:**

*CU\_AD\_FORMAT\_UNSIGNED\_INT8*

Unsigned 8-bit integers

*CU\_AD\_FORMAT\_UNSIGNED\_INT16*

Unsigned 16-bit integers

*CU\_AD\_FORMAT\_UNSIGNED\_INT32*

Unsigned 32-bit integers

*CU\_AD\_FORMAT\_SIGNED\_INT8*

Signed 8-bit integers

*CU\_AD\_FORMAT\_SIGNED\_INT16*

Signed 16-bit integers

*CU\_AD\_FORMAT\_SIGNED\_INT32*

Signed 32-bit integers

*CU\_AD\_FORMAT\_HALF*

16-bit floating point

*CU\_AD\_FORMAT\_FLOAT*

32-bit floating point



**enum CUcomputemode**

Compute Modes

**Enumerator:***CU\_COMPUTEMODE\_DEFAULT*

Default compute mode (Multiple contexts allowed per device)

*CU\_COMPUTEMODE\_EXCLUSIVE*

Compute-exclusive-thread mode (Only one context used by a single thread can be present on this device at a time)

*CU\_COMPUTEMODE\_PROHIBITED*

Compute-prohibited mode (No contexts can be created on this device at this time)

*CU\_COMPUTEMODE\_EXCLUSIVE\_PROCESS*

Compute-exclusive-process mode (Only one context used by a single process can be present on this device at a time)

**enum CUctx\_flags**

Context creation flags

**Enumerator:***CU\_CTX\_SCHED\_AUTO*

Automatic scheduling

*CU\_CTX\_SCHED\_SPIN*

Set spin as default scheduling

*CU\_CTX\_SCHED\_YIELD*

Set yield as default scheduling

*CU\_CTX\_SCHED\_BLOCKING\_SYNC*

Set blocking synchronization as default scheduling

*CU\_CTX\_BLOCKING\_SYNC*

Set blocking synchronization as default scheduling

**Deprecated**This flag was deprecated as of CUDA 4.0 and was replaced with **CU\_CTX\_SCHED\_BLOCKING\_SYNC**.*CU\_CTX\_MAP\_HOST*

Support mapped pinned allocations

*CU\_CTX\_LMEM\_RESIZE\_TO\_MAX*

Keep local memory allocation after launch

**enum CUdevice\_attribute**

Device properties

**Enumerator:***CU\_DEVICE\_ATTRIBUTE\_MAX\_THREADS\_PER\_BLOCK*

Maximum number of threads per block

*CU\_DEVICE\_ATTRIBUTE\_MAX\_BLOCK\_DIM\_X*

Maximum block dimension X

*CU\_DEVICE\_ATTRIBUTE\_MAX\_BLOCK\_DIM\_Y*

Maximum block dimension Y

*CU\_DEVICE\_ATTRIBUTE\_MAX\_BLOCK\_DIM\_Z*

Maximum block dimension Z

*CU\_DEVICE\_ATTRIBUTE\_MAX\_GRID\_DIM\_X*

Maximum grid dimension X

*CU\_DEVICE\_ATTRIBUTE\_MAX\_GRID\_DIM\_Y*

Maximum grid dimension Y



*CU\_DEVICE\_ATTRIBUTE\_MAX\_GRID\_DIM\_Z*  
Maximum grid dimension Z

*CU\_DEVICE\_ATTRIBUTE\_MAX\_SHARED\_MEMORY\_PER\_BLOCK*  
Maximum shared memory available per block in bytes

*CU\_DEVICE\_ATTRIBUTE\_SHARED\_MEMORY\_PER\_BLOCK*  
Deprecated, use *CU\_DEVICE\_ATTRIBUTE\_MAX\_SHARED\_MEMORY\_PER\_BLOCK*

*CU\_DEVICE\_ATTRIBUTE\_TOTAL\_CONSTANT\_MEMORY*  
Memory available on device for `__constant__` variables in a CUDA C kernel in bytes

*CU\_DEVICE\_ATTRIBUTE\_WARP\_SIZE*  
Warp size in threads

*CU\_DEVICE\_ATTRIBUTE\_MAX\_PITCH*  
Maximum pitch in bytes allowed by memory copies

*CU\_DEVICE\_ATTRIBUTE\_MAX\_REGISTERS\_PER\_BLOCK*  
Maximum number of 32-bit registers available per block

*CU\_DEVICE\_ATTRIBUTE\_REGISTERS\_PER\_BLOCK*  
Deprecated, use *CU\_DEVICE\_ATTRIBUTE\_MAX\_REGISTERS\_PER\_BLOCK*

*CU\_DEVICE\_ATTRIBUTE\_CLOCK\_RATE*  
Typical clock frequency in kilohertz

*CU\_DEVICE\_ATTRIBUTE\_TEXTURE\_ALIGNMENT*  
Alignment requirement for textures

*CU\_DEVICE\_ATTRIBUTE\_GPU\_OVERLAP*  
Device can possibly copy memory and execute a kernel concurrently. Deprecated. Use instead *CU\_DEVICE\_ATTRIBUTE\_ASYNC\_ENGINE\_COUNT*.

*CU\_DEVICE\_ATTRIBUTE\_MULTIPROCESSOR\_COUNT*  
Number of multiprocessors on device

*CU\_DEVICE\_ATTRIBUTE\_KERNEL\_EXEC\_TIMEOUT*  
Specifies whether there is a run time limit on kernels

*CU\_DEVICE\_ATTRIBUTE\_INTEGRATED*  
Device is integrated with host memory

*CU\_DEVICE\_ATTRIBUTE\_CAN\_MAP\_HOST\_MEMORY*  
Device can map host memory into CUDA address space

*CU\_DEVICE\_ATTRIBUTE\_COMPUTE\_MODE*  
Compute mode (See **CUcomputemode** for details)

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE1D\_WIDTH*  
Maximum 1D texture width

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE2D\_WIDTH*  
Maximum 2D texture width

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE2D\_HEIGHT*  
Maximum 2D texture height

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE3D\_WIDTH*  
Maximum 3D texture width

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE3D\_HEIGHT*  
Maximum 3D texture height

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE3D\_DEPTH*  
Maximum 3D texture depth

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE2D\_LAYERED\_WIDTH*  
Maximum 2D layered texture width



*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE2D\_LAYERED\_HEIGHT*  
Maximum 2D layered texture height

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE2D\_LAYERED\_LAYERS*  
Maximum layers in a 2D layered texture

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE2D\_ARRAY\_WIDTH*  
Deprecated, use  
*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE2D\_LAYERED\_WIDTH*

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE2D\_ARRAY\_HEIGHT*  
Deprecated, use  
*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE2D\_LAYERED\_HEIGHT*

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE2D\_ARRAY\_NUMSLICES*  
Deprecated, use  
*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE2D\_LAYERED\_LAYERS*

*CU\_DEVICE\_ATTRIBUTE\_SURFACE\_ALIGNMENT*  
Alignment requirement for surfaces

*CU\_DEVICE\_ATTRIBUTE\_CONCURRENT\_KERNELS*  
Device can possibly execute multiple kernels concurrently

*CU\_DEVICE\_ATTRIBUTE\_ECC\_ENABLED*  
Device has ECC support enabled

*CU\_DEVICE\_ATTRIBUTE\_PCI\_BUS\_ID*  
PCI bus ID of the device

*CU\_DEVICE\_ATTRIBUTE\_PCI\_DEVICE\_ID*  
PCI device ID of the device

*CU\_DEVICE\_ATTRIBUTE\_TCC\_DRIVER*  
Device is using TCC driver model

*CU\_DEVICE\_ATTRIBUTE\_MEMORY\_CLOCK\_RATE*  
Peak memory clock frequency in kilohertz

*CU\_DEVICE\_ATTRIBUTE\_GLOBAL\_MEMORY\_BUS\_WIDTH*  
Global memory bus width in bits

*CU\_DEVICE\_ATTRIBUTE\_L2\_CACHE\_SIZE*  
Size of L2 cache in bytes

*CU\_DEVICE\_ATTRIBUTE\_MAX\_THREADS\_PER\_MULTIPROCESSOR*  
Maximum resident threads per multiprocessor

*CU\_DEVICE\_ATTRIBUTE\_ASYNC\_ENGINE\_COUNT*  
Number of asynchronous engines

*CU\_DEVICE\_ATTRIBUTE\_UNIFIED\_ADDRESSING*  
Device shares a unified address space with the host

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE1D\_LAYERED\_WIDTH*  
Maximum 1D layered texture width

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE1D\_LAYERED\_LAYERS*  
Maximum layers in a 1D layered texture

*CU\_DEVICE\_ATTRIBUTE\_CAN\_TEX2D\_GATHER*  
Deprecated, do not use.

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE2D\_GATHER\_WIDTH*  
Maximum 2D texture width if *CUDA\_ARRAY3D\_TEXTURE\_GATHER* is set

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE2D\_GATHER\_HEIGHT*  
Maximum 2D texture height if *CUDA\_ARRAY3D\_TEXTURE\_GATHER* is set

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE3D\_WIDTH\_ALTERNATE*  
Alternate maximum 3D texture width



*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE3D\_HEIGHT\_ALTERNATE*  
Alternate maximum 3D texture height

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE3D\_DEPTH\_ALTERNATE*  
Alternate maximum 3D texture depth

*CU\_DEVICE\_ATTRIBUTE\_PCI\_DOMAIN\_ID*  
PCI domain ID of the device

*CU\_DEVICE\_ATTRIBUTE\_TEXTURE\_PITCH\_ALIGNMENT*  
Pitch alignment requirement for textures

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURECUBEMAP\_WIDTH*  
Maximum cubemap texture width/height

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURECUBEMAP\_LAYERED\_WIDTH*  
Maximum cubemap layered texture width/height

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURECUBEMAP\_LAYERED\_LAYERS*  
Maximum layers in a cubemap layered texture

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_SURFACE1D\_WIDTH*  
Maximum 1D surface width

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_SURFACE2D\_WIDTH*  
Maximum 2D surface width

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_SURFACE2D\_HEIGHT*  
Maximum 2D surface height

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_SURFACE3D\_WIDTH*  
Maximum 3D surface width

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_SURFACE3D\_HEIGHT*  
Maximum 3D surface height

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_SURFACE3D\_DEPTH*  
Maximum 3D surface depth

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_SURFACE1D\_LAYERED\_WIDTH*  
Maximum 1D layered surface width

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_SURFACE1D\_LAYERED\_LAYERS*  
Maximum layers in a 1D layered surface

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_SURFACE2D\_LAYERED\_WIDTH*  
Maximum 2D layered surface width

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_SURFACE2D\_LAYERED\_HEIGHT*  
Maximum 2D layered surface height

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_SURFACE2D\_LAYERED\_LAYERS*  
Maximum layers in a 2D layered surface

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_SURFACECUBEMAP\_WIDTH*  
Maximum cubemap surface width

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_SURFACECUBEMAP\_LAYERED\_WIDTH*  
Maximum cubemap layered surface width

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_SURFACECUBEMAP\_LAYERED\_LAYERS*  
Maximum layers in a cubemap layered surface

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE1D\_LINEAR\_WIDTH*  
Maximum 1D linear texture width

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE2D\_LINEAR\_WIDTH*  
Maximum 2D linear texture width

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE2D\_LINEAR\_HEIGHT*  
Maximum 2D linear texture height



*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE2D\_LINEAR\_PITCH*

Maximum 2D linear texture pitch in bytes

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE2D\_MIPMAPPED\_WIDTH*

Maximum mipmapped 2D texture width

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE2D\_MIPMAPPED\_HEIGHT*

Maximum mipmapped 2D texture height

*CU\_DEVICE\_ATTRIBUTE\_COMPUTE\_CAPABILITY\_MAJOR*

Major compute capability version number

*CU\_DEVICE\_ATTRIBUTE\_COMPUTE\_CAPABILITY\_MINOR*

Minor compute capability version number

*CU\_DEVICE\_ATTRIBUTE\_MAXIMUM\_TEXTURE1D\_MIPMAPPED\_WIDTH*

Maximum mipmapped 1D texture width

*CU\_DEVICE\_ATTRIBUTE\_STREAM\_PRIORITIES\_SUPPORTED*

Device supports stream priorities

*CU\_DEVICE\_ATTRIBUTE\_GLOBAL\_L1\_CACHE\_SUPPORTED*

Device supports caching globals in L1

*CU\_DEVICE\_ATTRIBUTE\_LOCAL\_L1\_CACHE\_SUPPORTED*

Device supports caching locals in L1

*CU\_DEVICE\_ATTRIBUTE\_MAX\_SHARED\_MEMORY\_PER\_MULTIPROCESSOR*

Maximum shared memory available per multiprocessor in bytes

*CU\_DEVICE\_ATTRIBUTE\_MAX\_REGISTERS\_PER\_MULTIPROCESSOR*

Maximum number of 32-bit registers available per multiprocessor

*CU\_DEVICE\_ATTRIBUTE\_MANAGED\_MEMORY*

Device can allocate managed memory on this system

*CU\_DEVICE\_ATTRIBUTE\_MULTI\_GPU\_BOARD*

Device is on a multi-GPU board

*CU\_DEVICE\_ATTRIBUTE\_MULTI\_GPU\_BOARD\_GROUP\_ID*

Unique id for a group of devices on the same multi-GPU board

#### **enum CUevent\_flags**

Event creation flags

##### **Enumerator:**

*CU\_EVENT\_DEFAULT*

Default event flag

*CU\_EVENT\_BLOCKING\_SYNC*

Event uses blocking synchronization

*CU\_EVENT\_DISABLE\_TIMING*

Event will not record timing data

*CU\_EVENT\_INTERPROCESS*

Event is suitable for interprocess use. *CU\_EVENT\_DISABLE\_TIMING* must be set

#### **enum CUfilter\_mode**

Texture reference filtering modes

##### **Enumerator:**

*CU\_TR\_FILTER\_MODE\_POINT*

Point filter mode

*CU\_TR\_FILTER\_MODE\_LINEAR*

Linear filter mode

#### **enum CUfunc\_cache**

Function cache configurations

##### **Enumerator:**



*CU\_FUNC\_CACHE\_PREFER\_NONE*

no preference for shared memory or L1 (default)

*CU\_FUNC\_CACHE\_PREFER\_SHARED*

prefer larger shared memory and smaller L1 cache

*CU\_FUNC\_CACHE\_PREFER\_L1*

prefer larger L1 cache and smaller shared memory

*CU\_FUNC\_CACHE\_PREFER\_EQUAL*

prefer equal sized L1 cache and shared memory

### **enum CUfunction\_attribute**

Function properties

#### **Enumerator:**

*CU\_FUNC\_ATTRIBUTE\_MAX\_THREADS\_PER\_BLOCK*

The maximum number of threads per block, beyond which a launch of the function would fail. This number depends on both the function and the device on which the function is currently loaded.

*CU\_FUNC\_ATTRIBUTE\_SHARED\_SIZE\_BYTES*

The size in bytes of statically-allocated shared memory required by this function. This does not include dynamically-allocated shared memory requested by the user at runtime.

*CU\_FUNC\_ATTRIBUTE\_CONST\_SIZE\_BYTES*

The size in bytes of user-allocated constant memory required by this function.

*CU\_FUNC\_ATTRIBUTE\_LOCAL\_SIZE\_BYTES*

The size in bytes of local memory used by each thread of this function.

*CU\_FUNC\_ATTRIBUTE\_NUM\_REGS*

The number of registers used by each thread of this function.

*CU\_FUNC\_ATTRIBUTE\_PTX\_VERSION*

The PTX virtual architecture version for which the function was compiled. This value is the major PTX version \* 10 + the minor PTX version, so a PTX version 1.3 function would return the value 13. Note that this may return the undefined value of 0 for cubins compiled prior to CUDA 3.0.

*CU\_FUNC\_ATTRIBUTE\_BINARY\_VERSION*

The binary architecture version for which the function was compiled. This value is the major binary version \* 10 + the minor binary version, so a binary version 1.3 function would return the value 13. Note that this will return a value of 10 for legacy cubins that do not have a properly-encoded binary architecture version.

*CU\_FUNC\_ATTRIBUTE\_CACHE\_MODE\_CA*

The attribute to indicate whether the function has been compiled with user specified option '-Xptxas --dlcm=ca' set .

### **enum CUgraphicsMapResourceFlags**

Flags for mapping and unmapping interop resources

### **enum CUgraphicsRegisterFlags**

Flags to register a graphics resource

### **enum CUipcMem\_flags**

CUDA Ipc Mem Flags

#### **Enumerator:**

*CU\_IPC\_MEM\_LAZY\_ENABLE\_PEER\_ACCESS*

Automatically enable peer access between remote devices as needed

### **enum CUjit\_cacheMode**

Caching modes for dlcm

#### **Enumerator:**



*CU\_JIT\_CACHE\_OPTION\_NONE*

Compile with no `-dlcm` flag specified

*CU\_JIT\_CACHE\_OPTION\_CG*

Compile with L1 cache disabled

*CU\_JIT\_CACHE\_OPTION\_CA*

Compile with L1 cache enabled

### **enum CUjit\_fallback**

Cubin matching fallback strategies

#### **Enumerator:**

*CU\_PREFER\_PTX*

Prefer to compile ptx if exact binary match not found

*CU\_PREFER\_BINARY*

Prefer to fall back to compatible binary code if exact match not found

### **enum CUjit\_option**

Online compiler and linker options

#### **Enumerator:**

*CU\_JIT\_MAX\_REGISTERS*

Max number of registers that a thread may use.

Option type: unsigned int

Applies to: compiler only

*CU\_JIT\_THREADS\_PER\_BLOCK*

IN: Specifies minimum number of threads per block to target compilation for

OUT: Returns the number of threads the compiler actually targeted. This restricts the resource utilization to the compiler (e.g. max registers) such that a block with the given number of threads should be able to launch based on register limitations. Note, this option does not currently take into account any other resource limitations, such as shared memory utilization.

Cannot be combined with **CU\_JIT\_TARGET**.

Option type: unsigned int

Applies to: compiler only

*CU\_JIT\_WALL\_TIME*

Overwrites the option value with the total wall clock time, in milliseconds, spent in the compiler and linker

Option type: float

Applies to: compiler and linker

*CU\_JIT\_INFO\_LOG\_BUFFER*

Pointer to a buffer in which to print any log messages that are informational in nature (the buffer size is specified via option **CU\_JIT\_INFO\_LOG\_BUFFER\_SIZE\_BYTES**)

Option type: char \*

Applies to: compiler and linker

*CU\_JIT\_INFO\_LOG\_BUFFER\_SIZE\_BYTES*

IN: Log buffer size in bytes. Log messages will be capped at this size (including null terminator)

OUT: Amount of log buffer filled with messages

Option type: unsigned int

Applies to: compiler and linker

*CU\_JIT\_ERROR\_LOG\_BUFFER*

Pointer to a buffer in which to print any log messages that reflect errors (the buffer size is specified via option **CU\_JIT\_ERROR\_LOG\_BUFFER\_SIZE\_BYTES**)

Option type: char \*

Applies to: compiler and linker

*CU\_JIT\_ERROR\_LOG\_BUFFER\_SIZE\_BYTES*

IN: Log buffer size in bytes. Log messages will be capped at this size (including null terminator)



OUT: Amount of log buffer filled with messages

Option type: unsigned int

Applies to: compiler and linker

#### *CU\_JIT\_OPTIMIZATION\_LEVEL*

Level of optimizations to apply to generated code (0 - 4), with 4 being the default and highest level of optimizations.

Option type: unsigned int

Applies to: compiler only

#### *CU\_JIT\_TARGET\_FROM\_CUCONTEXT*

No option value required. Determines the target based on the current attached context (default)

Option type: No option value needed

Applies to: compiler and linker

#### *CU\_JIT\_TARGET*

Target is chosen based on supplied **CUjit\_target**. Cannot be combined with

**CU\_JIT\_THREADS\_PER\_BLOCK**.

Option type: unsigned int for enumerated type **CUjit\_target**

Applies to: compiler and linker

#### *CU\_JIT\_FALLBACK\_STRATEGY*

Specifies choice of fallback strategy if matching cubin is not found. Choice is based on supplied **CUjit\_fallback**.

Option type: unsigned int for enumerated type **CUjit\_fallback**

Applies to: compiler only

#### *CU\_JIT\_GENERATE\_DEBUG\_INFO*

Specifies whether to create debug information in output (-g) (0: false, default)

Option type: int

Applies to: compiler and linker

#### *CU\_JIT\_LOG\_VERBOSE*

Generate verbose log messages (0: false, default)

Option type: int

Applies to: compiler and linker

#### *CU\_JIT\_GENERATE\_LINE\_INFO*

Generate line number information (-lineinfo) (0: false, default)

Option type: int

Applies to: compiler only

#### *CU\_JIT\_CACHE\_MODE*

Specifies whether to enable caching explicitly (-dlcm)

Choice is based on supplied **CUjit\_cacheMode\_enum**.

Option type: unsigned int for enumerated type **CUjit\_cacheMode\_enum**

Applies to: compiler only

### **enum CUjit\_target**

Online compilation targets

#### **Enumerator:**

##### *CU\_TARGET\_COMPUTE\_10*

Compute device class 1.0

##### *CU\_TARGET\_COMPUTE\_11*

Compute device class 1.1

##### *CU\_TARGET\_COMPUTE\_12*

Compute device class 1.2

##### *CU\_TARGET\_COMPUTE\_13*

Compute device class 1.3

##### *CU\_TARGET\_COMPUTE\_20*

Compute device class 2.0



*CU\_TARGET\_COMPUTE\_21*  
Compute device class 2.1

*CU\_TARGET\_COMPUTE\_30*  
Compute device class 3.0

*CU\_TARGET\_COMPUTE\_32*  
Compute device class 3.2

*CU\_TARGET\_COMPUTE\_35*  
Compute device class 3.5

*CU\_TARGET\_COMPUTE\_50*  
Compute device class 5.0

#### **enum CUjitInputType**

Device code formats

##### **Enumerator:**

*CU\_JIT\_INPUT\_CUBIN*  
Compiled device-class-specific device code  
Applicable options: none

*CU\_JIT\_INPUT\_PTX*  
PTX source code  
Applicable options: PTX compiler options

*CU\_JIT\_INPUT\_FATBINARY*  
Bundle of multiple cubins and/or PTX of some device code  
Applicable options: PTX compiler options, **CU\_JIT\_FALLBACK\_STRATEGY**

*CU\_JIT\_INPUT\_OBJECT*  
Host object with embedded device code  
Applicable options: PTX compiler options, **CU\_JIT\_FALLBACK\_STRATEGY**

*CU\_JIT\_INPUT\_LIBRARY*  
Archive of host objects with embedded device code  
Applicable options: PTX compiler options, **CU\_JIT\_FALLBACK\_STRATEGY**

#### **enum CULimit**

Limits

##### **Enumerator:**

*CU\_LIMIT\_STACK\_SIZE*  
GPU thread stack size

*CU\_LIMIT\_PRINTF\_FIFO\_SIZE*  
GPU printf FIFO size

*CU\_LIMIT\_MALLOC\_HEAP\_SIZE*  
GPU malloc heap size

*CU\_LIMIT\_DEV\_RUNTIME\_SYNC\_DEPTH*  
GPU device runtime launch synchronize depth

*CU\_LIMIT\_DEV\_RUNTIME\_PENDING\_LAUNCH\_COUNT*  
GPU device runtime pending launch count

#### **enum CUMemAttach\_flags**

CUDA Mem Attach Flags

##### **Enumerator:**

*CU\_MEM\_ATTACH\_GLOBAL*  
Memory can be accessed by any stream on any device

*CU\_MEM\_ATTACH\_HOST*  
Memory cannot be accessed by any stream on any device



*CU\_MEM\_ATTACH\_SINGLE*

Memory can only be accessed by a single stream on the associated device

### **enum CUmemorytype**

Memory types

#### **Enumerator:**

*CU\_MEMORYTYPE\_HOST*

Host memory

*CU\_MEMORYTYPE\_DEVICE*

Device memory

*CU\_MEMORYTYPE\_ARRAY*

Array memory

*CU\_MEMORYTYPE\_UNIFIED*

Unified device or host memory

### **enum CUpointer\_attribute**

Pointer information

#### **Enumerator:**

*CU\_POINTER\_ATTRIBUTE\_CONTEXT*

The **CUcontext** on which a pointer was allocated or registered

*CU\_POINTER\_ATTRIBUTE\_MEMORY\_TYPE*

The **CUmemorytype** describing the physical location of a pointer

*CU\_POINTER\_ATTRIBUTE\_DEVICE\_POINTER*

The address at which a pointer's memory may be accessed on the device

*CU\_POINTER\_ATTRIBUTE\_HOST\_POINTER*

The address at which a pointer's memory may be accessed on the host

*CU\_POINTER\_ATTRIBUTE\_P2P\_TOKENS*

A pair of tokens for use with the nv-p2p.h Linux kernel interface

*CU\_POINTER\_ATTRIBUTE\_SYNC\_MEMOPS*

Synchronize every synchronous memory operation initiated on this region

*CU\_POINTER\_ATTRIBUTE\_BUFFER\_ID*

A process-wide unique ID for an allocated memory region

*CU\_POINTER\_ATTRIBUTE\_IS\_MANAGED*

Indicates if the pointer points to managed memory

### **enum CUresourcetype**

Resource types

#### **Enumerator:**

*CU\_RESOURCE\_TYPE\_ARRAY*

Array resource

*CU\_RESOURCE\_TYPE\_MIPMAPPED\_ARRAY*

Mipmapped array resource

*CU\_RESOURCE\_TYPE\_LINEAR*

Linear resource

*CU\_RESOURCE\_TYPE\_PITCH2D*

Pitch 2D resource

### **enum CUresourceViewFormat**

Resource view format

#### **Enumerator:**

*CU\_RES\_VIEW\_FORMAT\_NONE*

No resource view format (use underlying resource format)



*CU\_RES\_VIEW\_FORMAT\_UINT\_1X8*  
 1 channel unsigned 8-bit integers  
*CU\_RES\_VIEW\_FORMAT\_UINT\_2X8*  
 2 channel unsigned 8-bit integers  
*CU\_RES\_VIEW\_FORMAT\_UINT\_4X8*  
 4 channel unsigned 8-bit integers  
*CU\_RES\_VIEW\_FORMAT\_SINT\_1X8*  
 1 channel signed 8-bit integers  
*CU\_RES\_VIEW\_FORMAT\_SINT\_2X8*  
 2 channel signed 8-bit integers  
*CU\_RES\_VIEW\_FORMAT\_SINT\_4X8*  
 4 channel signed 8-bit integers  
*CU\_RES\_VIEW\_FORMAT\_UINT\_1X16*  
 1 channel unsigned 16-bit integers  
*CU\_RES\_VIEW\_FORMAT\_UINT\_2X16*  
 2 channel unsigned 16-bit integers  
*CU\_RES\_VIEW\_FORMAT\_UINT\_4X16*  
 4 channel unsigned 16-bit integers  
*CU\_RES\_VIEW\_FORMAT\_SINT\_1X16*  
 1 channel signed 16-bit integers  
*CU\_RES\_VIEW\_FORMAT\_SINT\_2X16*  
 2 channel signed 16-bit integers  
*CU\_RES\_VIEW\_FORMAT\_SINT\_4X16*  
 4 channel signed 16-bit integers  
*CU\_RES\_VIEW\_FORMAT\_UINT\_1X32*  
 1 channel unsigned 32-bit integers  
*CU\_RES\_VIEW\_FORMAT\_UINT\_2X32*  
 2 channel unsigned 32-bit integers  
*CU\_RES\_VIEW\_FORMAT\_UINT\_4X32*  
 4 channel unsigned 32-bit integers  
*CU\_RES\_VIEW\_FORMAT\_SINT\_1X32*  
 1 channel signed 32-bit integers  
*CU\_RES\_VIEW\_FORMAT\_SINT\_2X32*  
 2 channel signed 32-bit integers  
*CU\_RES\_VIEW\_FORMAT\_SINT\_4X32*  
 4 channel signed 32-bit integers  
*CU\_RES\_VIEW\_FORMAT\_FLOAT\_1X16*  
 1 channel 16-bit floating point  
*CU\_RES\_VIEW\_FORMAT\_FLOAT\_2X16*  
 2 channel 16-bit floating point  
*CU\_RES\_VIEW\_FORMAT\_FLOAT\_4X16*  
 4 channel 16-bit floating point  
*CU\_RES\_VIEW\_FORMAT\_FLOAT\_1X32*  
 1 channel 32-bit floating point  
*CU\_RES\_VIEW\_FORMAT\_FLOAT\_2X32*  
 2 channel 32-bit floating point  
*CU\_RES\_VIEW\_FORMAT\_FLOAT\_4X32*  
 4 channel 32-bit floating point



*CU\_RES\_VIEW\_FORMAT\_UNSIGNED\_BC1*  
Block compressed 1

*CU\_RES\_VIEW\_FORMAT\_UNSIGNED\_BC2*  
Block compressed 2

*CU\_RES\_VIEW\_FORMAT\_UNSIGNED\_BC3*  
Block compressed 3

*CU\_RES\_VIEW\_FORMAT\_UNSIGNED\_BC4*  
Block compressed 4 unsigned

*CU\_RES\_VIEW\_FORMAT\_SIGNED\_BC4*  
Block compressed 4 signed

*CU\_RES\_VIEW\_FORMAT\_UNSIGNED\_BC5*  
Block compressed 5 unsigned

*CU\_RES\_VIEW\_FORMAT\_SIGNED\_BC5*  
Block compressed 5 signed

*CU\_RES\_VIEW\_FORMAT\_UNSIGNED\_BC6H*  
Block compressed 6 unsigned half-float

*CU\_RES\_VIEW\_FORMAT\_SIGNED\_BC6H*  
Block compressed 6 signed half-float

*CU\_RES\_VIEW\_FORMAT\_UNSIGNED\_BC7*  
Block compressed 7

**enum CUresult**

Error codes

**Enumerator:***CUDA\_SUCCESS*

The API call returned with no errors. In the case of query calls, this can also mean that the operation being queried is complete (see **cuEventQuery()** and **cuStreamQuery()**).

*CUDA\_ERROR\_INVALID\_VALUE*

This indicates that one or more of the parameters passed to the API call is not within an acceptable range of values.

*CUDA\_ERROR\_OUT\_OF\_MEMORY*

The API call failed because it was unable to allocate enough memory to perform the requested operation.

*CUDA\_ERROR\_NOT\_INITIALIZED*

This indicates that the CUDA driver has not been initialized with **cuInit()** or that initialization has failed.

*CUDA\_ERROR\_DEINITIALIZED*

This indicates that the CUDA driver is in the process of shutting down.

*CUDA\_ERROR\_PROFILER\_DISABLED*

This indicates profiler is not initialized for this run. This can happen when the application is running with external profiling tools like visual profiler.

*CUDA\_ERROR\_PROFILER\_NOT\_INITIALIZED***Deprecated**

This error return is deprecated as of CUDA 5.0. It is no longer an error to attempt to enable/disable the profiling via **cuProfilerStart** or **cuProfilerStop** without initialization.

*CUDA\_ERROR\_PROFILER\_ALREADY\_STARTED***Deprecated**

This error return is deprecated as of CUDA 5.0. It is no longer an error to call **cuProfilerStart()** when profiling is already enabled.

*CUDA\_ERROR\_PROFILER\_ALREADY\_STOPPED***Deprecated**

This error return is deprecated as of CUDA 5.0. It is no longer an error to call **cuProfilerStop()**



when profiling is already disabled.

#### *CUDA\_ERROR\_NO\_DEVICE*

This indicates that no CUDA-capable devices were detected by the installed CUDA driver.

#### *CUDA\_ERROR\_INVALID\_DEVICE*

This indicates that the device ordinal supplied by the user does not correspond to a valid CUDA device.

#### *CUDA\_ERROR\_INVALID\_IMAGE*

This indicates that the device kernel image is invalid. This can also indicate an invalid CUDA module.

#### *CUDA\_ERROR\_INVALID\_CONTEXT*

This most frequently indicates that there is no context bound to the current thread. This can also be returned if the context passed to an API call is not a valid handle (such as a context that has had **cuCtxDestroy()** invoked on it). This can also be returned if a user mixes different API versions (i.e. 3010 context with 3020 API calls). See **cuCtxGetApiVersion()** for more details.

#### *CUDA\_ERROR\_CONTEXT\_ALREADY\_CURRENT*

This indicated that the context being supplied as a parameter to the API call was already the active context.

#### **Deprecated**

This error return is deprecated as of CUDA 3.2. It is no longer an error to attempt to push the active context via **cuCtxPushCurrent()**.

#### *CUDA\_ERROR\_MAP\_FAILED*

This indicates that a map or register operation has failed.

#### *CUDA\_ERROR\_UNMAP\_FAILED*

This indicates that an unmap or unregister operation has failed.

#### *CUDA\_ERROR\_ARRAY\_IS\_MAPPED*

This indicates that the specified array is currently mapped and thus cannot be destroyed.

#### *CUDA\_ERROR\_ALREADY\_MAPPED*

This indicates that the resource is already mapped.

#### *CUDA\_ERROR\_NO\_BINARY\_FOR\_GPU*

This indicates that there is no kernel image available that is suitable for the device. This can occur when a user specifies code generation options for a particular CUDA source file that do not include the corresponding device configuration.

#### *CUDA\_ERROR\_ALREADY\_ACQUIRED*

This indicates that a resource has already been acquired.

#### *CUDA\_ERROR\_NOT\_MAPPED*

This indicates that a resource is not mapped.

#### *CUDA\_ERROR\_NOT\_MAPPED\_AS\_ARRAY*

This indicates that a mapped resource is not available for access as an array.

#### *CUDA\_ERROR\_NOT\_MAPPED\_AS\_POINTER*

This indicates that a mapped resource is not available for access as a pointer.

#### *CUDA\_ERROR\_ECC\_UNCORRECTABLE*

This indicates that an uncorrectable ECC error was detected during execution.

#### *CUDA\_ERROR\_UNSUPPORTED\_LIMIT*

This indicates that the **CUlimit** passed to the API call is not supported by the active device.

#### *CUDA\_ERROR\_CONTEXT\_ALREADY\_IN\_USE*

This indicates that the **CUcontext** passed to the API call can only be bound to a single CPU thread at a time but is already bound to a CPU thread.

#### *CUDA\_ERROR\_PEER\_ACCESS\_UNSUPPORTED*

This indicates that peer access is not supported across the given devices.



**CUDA\_ERROR\_INVALID\_PTX**

This indicates that a PTX JIT compilation failed.

**CUDA\_ERROR\_INVALID\_SOURCE**

This indicates that the device kernel source is invalid.

**CUDA\_ERROR\_FILE\_NOT\_FOUND**

This indicates that the file specified was not found.

**CUDA\_ERROR\_SHARED\_OBJECT\_SYMBOL\_NOT\_FOUND**

This indicates that a link to a shared object failed to resolve.

**CUDA\_ERROR\_SHARED\_OBJECT\_INIT\_FAILED**

This indicates that initialization of a shared object failed.

**CUDA\_ERROR\_OPERATING\_SYSTEM**

This indicates that an OS call failed.

**CUDA\_ERROR\_INVALID\_HANDLE**

This indicates that a resource handle passed to the API call was not valid. Resource handles are opaque types like **CUstream** and **CUevent**.

**CUDA\_ERROR\_NOT\_FOUND**

This indicates that a named symbol was not found. Examples of symbols are global/constant variable names, texture names, and surface names.

**CUDA\_ERROR\_NOT\_READY**

This indicates that asynchronous operations issued previously have not completed yet. This result is not actually an error, but must be indicated differently than **CUDA\_SUCCESS** (which indicates completion). Calls that may return this value include **cuEventQuery()** and **cuStreamQuery()**.

**CUDA\_ERROR\_ILLEGAL\_ADDRESS**

While executing a kernel, the device encountered a load or store instruction on an invalid memory address. The context cannot be used, so it must be destroyed (and a new one should be created). All existing device memory allocations from this context are invalid and must be reconstructed if the program is to continue using CUDA.

**CUDA\_ERROR\_LAUNCH\_OUT\_OF\_RESOURCES**

This indicates that a launch did not occur because it did not have appropriate resources. This error usually indicates that the user has attempted to pass too many arguments to the device kernel, or the kernel launch specifies too many threads for the kernel's register count. Passing arguments of the wrong size (i.e. a 64-bit pointer when a 32-bit int is expected) is equivalent to passing too many arguments and can also result in this error.

**CUDA\_ERROR\_LAUNCH\_TIMEOUT**

This indicates that the device kernel took too long to execute. This can only occur if timeouts are enabled - see the device attribute

**CU\_DEVICE\_ATTRIBUTE\_KERNEL\_EXEC\_TIMEOUT** for more information. The context cannot be used (and must be destroyed similar to

**CUDA\_ERROR\_LAUNCH\_FAILED**). All existing device memory allocations from this context are invalid and must be reconstructed if the program is to continue using CUDA.

**CUDA\_ERROR\_LAUNCH\_INCOMPATIBLE\_TEXTURING**

This error indicates a kernel launch that uses an incompatible texturing mode.

**CUDA\_ERROR\_PEER\_ACCESS\_ALREADY\_ENABLED**

This error indicates that a call to **cuCtxEnablePeerAccess()** is trying to re-enable peer access to a context which has already had peer access to it enabled.

**CUDA\_ERROR\_PEER\_ACCESS\_NOT\_ENABLED**

This error indicates that **cuCtxDisablePeerAccess()** is trying to disable peer access which has not been enabled yet via **cuCtxEnablePeerAccess()**.

**CUDA\_ERROR\_PRIMARY\_CONTEXT\_ACTIVE**

This error indicates that the primary context for the specified device has already been initialized.



*CUDA\_ERROR\_CONTEXT\_IS\_DESTROYED*

This error indicates that the context current to the calling thread has been destroyed using **cuCtxDestroy**, or is a primary context which has not yet been initialized.

*CUDA\_ERROR\_ASSERT*

A device-side assert triggered during kernel execution. The context cannot be used anymore, and must be destroyed. All existing device memory allocations from this context are invalid and must be reconstructed if the program is to continue using CUDA.

*CUDA\_ERROR\_TOO\_MANY\_PEERS*

This error indicates that the hardware resources required to enable peer access have been exhausted for one or more of the devices passed to **cuCtxEnablePeerAccess()**.

*CUDA\_ERROR\_HOST\_MEMORY\_ALREADY\_REGISTERED*

This error indicates that the memory range passed to **cuMemHostRegister()** has already been registered.

*CUDA\_ERROR\_HOST\_MEMORY\_NOT\_REGISTERED*

This error indicates that the pointer passed to **cuMemHostUnregister()** does not correspond to any currently registered memory region.

*CUDA\_ERROR\_HARDWARE\_STACK\_ERROR*

While executing a kernel, the device encountered a stack error. This can be due to stack corruption or exceeding the stack size limit. The context cannot be used, so it must be destroyed (and a new one should be created). All existing device memory allocations from this context are invalid and must be reconstructed if the program is to continue using CUDA.

*CUDA\_ERROR\_ILLEGAL\_INSTRUCTION*

While executing a kernel, the device encountered an illegal instruction. The context cannot be used, so it must be destroyed (and a new one should be created). All existing device memory allocations from this context are invalid and must be reconstructed if the program is to continue using CUDA.

*CUDA\_ERROR\_MISALIGNED\_ADDRESS*

While executing a kernel, the device encountered a load or store instruction on a memory address which is not aligned. The context cannot be used, so it must be destroyed (and a new one should be created). All existing device memory allocations from this context are invalid and must be reconstructed if the program is to continue using CUDA.

*CUDA\_ERROR\_INVALID\_ADDRESS\_SPACE*

While executing a kernel, the device encountered an instruction which can only operate on memory locations in certain address spaces (global, shared, or local), but was supplied a memory address not belonging to an allowed address space. The context cannot be used, so it must be destroyed (and a new one should be created). All existing device memory allocations from this context are invalid and must be reconstructed if the program is to continue using CUDA.

*CUDA\_ERROR\_INVALID\_PC*

While executing a kernel, the device program counter wrapped its address space. The context cannot be used, so it must be destroyed (and a new one should be created). All existing device memory allocations from this context are invalid and must be reconstructed if the program is to continue using CUDA.

*CUDA\_ERROR\_LAUNCH\_FAILED*

An exception occurred on the device while executing a kernel. Common causes include dereferencing an invalid device pointer and accessing out of bounds shared memory. The context cannot be used, so it must be destroyed (and a new one should be created). All existing device memory allocations from this context are invalid and must be reconstructed if the program is to continue using CUDA.

*CUDA\_ERROR\_NOT\_PERMITTED*

This error indicates that the attempted operation is not permitted.

*CUDA\_ERROR\_NOT\_SUPPORTED*

This error indicates that the attempted operation is not supported on the current system or device.



*CUDA\_ERROR\_UNKNOWN*

This indicates that an unknown internal error has occurred.

**enum CUsharedconfig**

Shared memory configurations

**Enumerator:**

*CU\_SHARED\_MEM\_CONFIG\_DEFAULT\_BANK\_SIZE*

set default shared memory bank size

*CU\_SHARED\_MEM\_CONFIG\_FOUR\_BYTE\_BANK\_SIZE*

set shared memory bank width to four bytes

*CU\_SHARED\_MEM\_CONFIG\_EIGHT\_BYTE\_BANK\_SIZE*

set shared memory bank width to eight bytes

**enum CUstream\_flags**

Stream creation flags

**Enumerator:**

*CU\_STREAM\_DEFAULT*

Default stream flag

*CU\_STREAM\_NON\_BLOCKING*

Stream does not synchronize with stream 0 (the NULL stream)

**Author**

Generated automatically by Doxygen from the source code.

