



NVIDIA TensorRT

API Migration Guide | NVIDIA Docs

Table of Contents

Chapter 1. Python.....	1
1.1. Python API Changes.....	1
1.2. Added Python APIs.....	3
1.3. Removed Python APIs.....	4
Chapter 2. C++.....	7
2.1. C++ API Changes.....	7
2.2. 64-Bit Dimension Changes.....	8
2.3. Added C++ APIs.....	8
2.4. Removed C++ APIs.....	10
2.5. Removed C++ Plugins.....	14
2.6. Removed Safety C++ APIs.....	15
Chapter 3. trtexec.....	16
3.1. trtexec Flag Changes.....	16
3.2. Removed trtexec Flags.....	16
3.3. Deprecated trtexec Flags.....	16

Chapter 1. Python

1.1. Python API Changes

Table 1. Allocating Buffers and Using a Name-Based Engine API

TensorRT 8.x	TensorRT 10.0
<pre>def allocate_buffers(self, engine): """ Allocates all buffers required for an engine, i.e. host/device inputs/ outputs. """ inputs = [] outputs = [] bindings = [] stream = cuda.Stream() # binding is the name of input/ output for binding in engine: size = trt.volume(engine.get_binding_shape(binding) * engine.max_batch_size dtype = trt.nptype(engine.get_binding_dtype(binding)) # Allocate host and device buffers host_mem = cuda.pagelocked_empty(size, dtype) # page-locked memory buffer (won't swapped to disk) device_mem = cuda.mem_alloc(host_mem.nbytes) # Append the device buffer address to device bindings. # When cast to int, it's a linear index into the context's memory (like memory address). bindings.append(int(device_mem))</pre>	<pre>def allocate_buffers(self, engine): """ Allocates all buffers required for an engine, i.e. host/device inputs/ outputs. """ inputs = [] outputs = [] bindings = [] stream = cuda.Stream() for i in range(engine.num_io_tensors): tensor_name = engine.get_tensor_name(i) size = trt.volume(engine.get_tensor_shape(tensor_name)) dtype = trt.nptype(engine.get_tensor_dtype(tensor_name)) # Allocate host and device buffers host_mem = cuda.pagelocked_empty(size, dtype) # page-locked memory buffer (won't swapped to disk) device_mem = cuda.mem_alloc(host_mem.nbytes) # Append the device buffer address to device bindings. # When cast to int, it's a linear index into the context's memory (like memory address). bindings.append(int(device_mem)) # Append to the appropriate input/output list.</pre>

TensorRT 8.x	TensorRT 10.0
<pre># Append to the appropriate input/output list. if engine.binding_is_input(binding): inputs.append(self.HostDeviceMem(host_mem, device_mem)) else: outputs.append(self.HostDeviceMem(host_mem, device_mem)) return inputs, outputs, bindings, stream</pre>	<pre>if engine.get_tensor_mode(tensor_name) == trt.TensorIOMode.INPUT: inputs.append(self.HostDeviceMem(host_mem, device_mem)) else: outputs.append(self.HostDeviceMem(host_mem, device_mem)) return inputs, outputs, bindings, stream</pre>

Table 2. Transition from enqueueV2 to enqueueV3 for Python

TensorRT 8.x	TensorRT 10.0
<pre># Allocate device memory for inputs. d_inputs = [cuda.mem_alloc(input_nbytes) for binding in range(input_num)] # Allocate device memory for outputs. h_output = cuda.pagelocked_empty(output_nbytes, dtype=np.float32) d_output = cuda.mem_alloc(h_output.nbytes) # Transfer data from host to device. cuda.memcpy_htod_async(d_inputs[0], input_a, stream) cuda.memcpy_htod_async(d_inputs[1], input_b, stream) cuda.memcpy_htod_async(d_inputs[2], input_c, stream) # Run inference context.execute_async_v2(bindings=[int(d_inputs[i]) for i in range(3)] + [int(d_output)]), stream_handle=stream.handle) # Synchronize the stream stream.synchronize()</pre>	<pre># Allocate device memory for inputs. d_inputs = [cuda.mem_alloc(input_nbytes) for binding in range(input_num)] # Allocate device memory for outputs. h_output = cuda.pagelocked_empty(output_nbytes, dtype=np.float32) d_output = cuda.mem_alloc(h_output.nbytes) # Transfer data from host to device. cuda.memcpy_htod_async(d_inputs[0], input_a, stream) cuda.memcpy_htod_async(d_inputs[1], input_b, stream) cuda.memcpy_htod_async(d_inputs[2], input_c, stream) # Setup tensor address bindings = [int(d_inputs[i]) for i in range(3)] + [int(d_output)] for i in range(engine.num_io_tensors): context.set_tensor_address(engine.get_tensor_name(i), bindings[i]) # Run inference context.execute_async_v3(stream_handle=stream.handle) # Synchronize the stream stream.synchronize()</pre>

Table 3. Engine Building, use only build_serialized_network

TensorRT 8.x	TensorRT 10.0
<pre>engine_bytes = None try:</pre>	<pre>engine_bytes = self.builder.build_serialized_network(self.network, self.config)</pre>

TensorRT 8.x	TensorRT 10.0
<pre> engine_bytes = self.builder.build_serialized_network(self: self.config) except AttributeError: engine = self.builder.build_engine(self.network, self.config) engine_bytes = engine.serialize() del engine assert engine_bytes </pre>	<pre> if engine_bytes is None: log.error("Failed to create engine") sys.exit(1) </pre>

1.2. Added Python APIs

Types

- ▶ `APILanguage`
- ▶ `ExecutionContextAllocationStrategy`
- ▶ `IGpuAsyncAllocator`
- ▶ `InterfaceInfo`
- ▶ `IPluginResource`
- ▶ `IPluginV3`
- ▶ `IStreamReader`
- ▶ `IVersionedInterface`

Methods and Properties

- ▶ `ICudaEngine.is_debug_tensor()`
- ▶ `ICudaEngine.minimum_weight_streaming_budget`
- ▶ `ICudaEngine.streamable_weights_size`
- ▶ `ICudaEngine.weight_streaming_budget`
- ▶ `IEExecutionContext.get_debug_listener()`
- ▶ `IEExecutionContext.get_debug_state()`
- ▶ `IEExecutionContext.set_all_tensors_debug_state()`
- ▶ `IEExecutionContext.set_debug_listener()`
- ▶ `IEExecutionContext.set_tensor_debug_state()`
- ▶ `IEExecutionContext.update_device_memory_size_for_shapes()`
- ▶ `IGpuAllocator.allocate_async()`
- ▶ `IGpuAllocator.deallocate_async()`
- ▶ `INetworkDefinition.add_plugin_v3()`
- ▶ `INetworkDefinition.is_debug_tensor()`

- ▶ `INetworkDefinition.mark_debug()`
- ▶ `INetworkDefinition.unmark_debug()`
- ▶ `IPluginRegistry.acquire_plugin_resource()`
- ▶ `IPluginRegistry.all_creators`
- ▶ `IPluginRegistry.deregister_creator()`
- ▶ `IPluginRegistry.get_creator()`
- ▶ `IPluginRegistry.register_creator()`
- ▶ `IPluginRegistry.release_plugin_resource()`

1.3. Removed Python APIs

Table 4. Removed Python APIs and their Suggested Superseded API

Python API	Superseded API
<code>BuilderFlag.ENABLE_TACTIC_HEURISTIC</code>	Builder optimization level 2
<code>BuilderFlag.STRICT_TYPES</code>	Use all three flags: 1. <code>BuilderFlag.DIRECT_IO</code> 2. <code>BuilderFlag.PREFER_PRECISION_CONSTRAINTS</code> 3. <code>BuilderFlag.REJECT_EMPTY_ALGORITHMS</code>
1. <code>EngineCapability.DEFAULT</code> 2. <code>EngineCapability.kSAFE_DLA</code> 3. <code>EngineCapability.SAFE_GPU</code>	1. <code>EngineCapability.STANDARD</code> 2. <code>EngineCapability.DLA_STANDALONE</code> 3. <code>EngineCapability.SAFETY</code>
<code>IAlgorithmIOInfo.tensor_format</code>	The strides, data type, and vectorization information is sufficient to uniquely identify tensor formats.
<code>IBuilder.max_batch_size</code>	Implicit batch is no longer supported.
<code>IBuilderConfig.max_workspace_size</code>	1. <code>IBuilderConfig.set_memory_pool_limit()</code> with <code>MemoryPoolType.WORKSPACE</code> 2. <code>IBuilderConfig.get_memory_pool_limit()</code> with <code>MemoryPoolType.WORKSPACE</code>
<code>IBuilderConfig.min_timing_iterations</code>	<code>IBuilderConfig.avg_timing_iterations</code>
1. <code>ICudaEngine.binding_is_input()</code> 2. <code>ICudaEngine.get_binding_bytes_per_component()</code> 3. <code>ICudaEngine.get_binding_components_per_element()</code>	1. <code>ICudaEngine.get_tensor_mode()</code> 2. <code>ICudaEngine.get_tensor_bytes_per_component()</code> 3. <code>ICudaEngine.get_tensor_components_per_element()</code>

Python API	Superseded API
- ICudaEngine.get_binding_dtype() - ICudaEngine.get_binding_format() - ICudaEngine.get_binding_format_desc() - ICudaEngine.get_binding_index() - ICudaEngine.get_binding_name() - ICudaEngine.get_binding_shape() - ICudaEngine.get_binding_vectorized_dim() - ICudaEngine.get_location() - ICudaEngine.get_profile_shape() - ICudaEngine.get_profile_shape_input() - ICudaEngine.has_implicit_batch_dimension() - ICudaEngine.is_execution_binding() - ICudaEngine.is_shape_binding() - ICudaEngine.max_batch_size() - ICudaEngine.num_bindings()	- ICudaEngine.get_tensor_dtype() - ICudaEngine.get_tensor_format() - ICudaEngine.get_tensor_format_desc() - No name-based equivalent replacement - No name-based equivalent replacement - ICudaEngine.get_tensor_shape() - ICudaEngine.get_tensor_vectorized_dim() - ITensor.location - ICudaEngine.get_tensor_profile_shape() - ICudaEngine.get_tensor_profile_values() - Implicit batch is no longer supported - No name-based equivalent replacement - ICudaEngine.is_shape_inference_io() - Implicit batch is no longer supported - ICudaEngine.num_io_tensors()
- IExecutionContext.get_binding_shape() - IExecutionContext.get_strides() - IExecutionContext.set_binding_shape()	- IExecutionContext.get_tensor_shape() - IExecutionContext.get_tensor_strides() - IExecutionContext.set_input_shape()
IFullyConnectedLayer	IMatrixMultiplyLayer
- INetworkDefinition.add_convolution() - INetworkDefinition.add_deconvolution() - INetworkDefinition.add_fully_connected() - INetworkDefinition.add_padding() - INetworkDefinition.add_pooling() - INetworkDefinition.add_rnn_v2() - INetworkDefinition.has_explicit_precision - INetworkDefinition.has_implicit_batch_dimension	- INetworkDefinition.add_convolution_nd() - INetworkDefinition.add_deconvolution_nd() - INetworkDefinition.add_matrix_multiply() - INetworkDefinition.add_padding_nd() - INetworkDefinition.add_pooling_nd() - INetworkDefinition.add_loop() - Explicit precision support is removed in 10.0 - Implicit batch is no longer supported
IRNNv2Layer	ILoop
- NetworkDefinitionCreationFlag.EXPLICIT_BATCH - NetworkDefinitionCreationFlag.EXPLICIT_PRECISION	Support is removed in 10.0
PaddingMode.CAFFE_ROUND_DOWN	Caffe is not supported since 9.0

Python API	Superseded API
2. <code>PaddingMode.CAFFE_ROUND_UP</code>	
1. <code>PreviewFeature.DISABLE_EXTERNAL_TACTICS_0805</code> 2. <code>PreviewFeature.FASTER_DYNAMIC_SHAPES_0805</code>	1. External tactics are always disabled for core code 2. This flag is on by default
1. <code>ProfilingVerbosity.DEFAULT</code> 2. <code>ProfilingVerbosity.VERBOSE</code>	1. <code>ProfilingVerbosity.LAYER_NAMES_ONLY</code> 2. <code>ProfilingVerbosity.DETAILED</code>
<code>ResizeMode</code>	Use <code>InterpolationMode</code> , alias is removed
<code>SampleMode.DEFAULT</code>	<code>SampleMode.STRICT_BOUNDS</code>
<code>SliceMode</code>	Use <code>SampleMode</code> , alias is removed

Chapter 2. C++

2.1. C++ API Changes

Table 5. Transition from enqueueV2 to enqueueV3 for C++

TensorRT 8.x	TensorRT 10.0
<pre>// Create RAII buffer manager object. samplesCommon::BufferManager buffers(mEngine); auto context = SampleUniquePtr<nvinfer1::IExecutionContext> >createExecutionContext(); if (!context) { return false; } // Pick a random digit to try to infer. srand(time(NULL)); int32_t const digit = rand() % 10; // Read the input data into the managed buffers. // There should be just 1 input tensor. ASSERT(mParams.inputTensorNames.size() == 1); if (!processInput(buffers, mParams.inputTensorNames[0], digit)) { return false; } // Create CUDA stream for the execution of this inference. cudaStream_t stream; CHECK(cudaStreamCreate(&stream));</pre>	<pre>// Create RAII buffer manager object. samplesCommon::BufferManager buffers(mEngine); auto context = SampleUniquePtr<nvinfer1::IExecutionContext>(mEngine- >createExecutionContext()); if (!context) { return false; } for (int32_t i = 0, e = mEngine- >getNbIOTensors(); i < e; i++) { auto const name = mEngine- >getIOTensorName(i); context->setTensorAddress(name, buffers.getDeviceBuffer(name)); } // Pick a random digit to try to infer. srand(time(NULL)); int32_t const digit = rand() % 10; // Read the input data into the managed buffers. // There should be just 1 input tensor. ASSERT(mParams.inputTensorNames.size() == 1); if (!processInput(buffers, mParams.inputTensorNames[0], digit)) { return false; } // Create CUDA stream for the execution of this inference.</pre>

TensorRT 8.x	TensorRT 10.0
<pre>// Asynchronously copy data from host input buffers to device input buffers buffers.copyInputToDeviceAsync(stream); // Asynchronously enqueue the inference work if (!context->enqueueV2(buffers.getDeviceBindings().data, stream, nullptr)) { return false; } // Asynchronously copy data from device output buffers to host output buffers. buffers.copyOutputToHostAsync(stream); // Wait for the work in the stream to complete. CHECK(cudaStreamSynchronize(stream)); // Release stream. CHECK(cudaStreamDestroy(stream));</pre>	<pre>cudaStream_t stream; CHECK(cudaStreamCreate(&stream)); // Asynchronously copy data from host input buffers to device input buffers buffers.copyInputToDeviceAsync(stream); // Asynchronously enqueue the inference work if (!context->enqueueV3(stream)) { return false; } // Asynchronously copy data from device output buffers to host output buffers. buffers.copyOutputToHostAsync(stream); // Wait for the work in the stream to complete. CHECK(cudaStreamSynchronize(stream)); // Release stream. CHECK(cudaStreamDestroy(stream));</pre>

2.2. 64-Bit Dimension Changes

The dimensions held by Dims changed from `int32_t` to `int64_t`. However, in TensorRT 10.0, TensorRT will generally reject networks that use dimensions exceeding the range of `int32_t`. The tensor type returned by `IShapeLayer` is now `DataType::kINT64`. Use `ICastLayer` to cast the result to the tensor of type `DataType::kINT32` if 32-bit dimensions are required.

Inspect code that bitwise copies to and from Dims to ensure it is correct for `int64_t` dimensions.

2.3. Added C++ APIs

Enums

- ▶ `ActivationType::kGELU_ERF`
- ▶ `ActivationType::kGELU_TANH`
- ▶ `BuilderFlag::kREFIT_IDENTICAL`
- ▶ `BuilderFlag::kSTRIP_PLAN`
- ▶ `BuilderFlag::kWEIGHT_STREAMING`
- ▶ `Datatype::kINT4`
- ▶ `LayerType::kPLUGIN_V3`

Types

- ▶ `APILanguage`
- ▶ `Dims64`
- ▶ `ExecutionContextAllocationStrategy`
- ▶ `IGpuAsyncAllocator`
- ▶ `InterfaceInfo`
- ▶ `IPluginResource`
- ▶ `IPluginV3`
- ▶ `IStreamReader`
- ▶ `IVersionedInterface`

Methods and Properties

- ▶ `getInferLibBuildVersion`
- ▶ `getInferLibMajorVersion`
- ▶ `getInferLibMinorVersion`
- ▶ `getInferLibPatchVersion`
- ▶ `ICudaEngine::createRefitter`
- ▶ `IcudaEngine::getMinimumWeightStreamingBudget`
- ▶ `IcudaEngine::getStreamableWeightsSize`
- ▶ `ICudaEngine::getWeightStreamingBudget`
- ▶ `IcudaEngine::isDebugTensor`
- ▶ `ICudaEngine::setWeightStreamingBudget`
- ▶ `IExecutionContext::getDebugListener`
- ▶ `IExecutionContext::getTensorDebugState`
- ▶ `IExecutionContext::setAllTensorsDebugState`
- ▶ `IExecutionContext::setDebugListener`
- ▶ `IExecutionContext::setOutputTensorAddress`
- ▶ `IExecutionContext::setTensorDebugState`
- ▶ `IExecutionContext::updateDeviceMemorySizeForShapes`
- ▶ `IGpuAllocator::allocateAsync`
- ▶ `IGpuAllocator::deallocateAsync`
- ▶ `INetworkDefinition::addPluginV3`
- ▶ `INetworkDefinition::isDebugTensor`
- ▶ `INetworkDefinition::markDebug`
- ▶ `INetworkDefinition::unmarkDebug`

- ▶ `IPluginRegistry::acquirePluginResource`
- ▶ `IPluginRegistry::deregisterCreator`
- ▶ `IPluginRegistry::getAllCreators`
- ▶ `IPluginRegistry::getCreator`
- ▶ `IPluginRegistry::registerCreator`
- ▶ `IPluginRegistry::releasePluginResource`

2.4. Removed C++ APIs

Table 6. Removed C++ APIs and their Suggested Superseded API

C++ API	Superseded API
<code>BuilderFlag::kENABLE_TACTIC_HEURISTIC</code>	Builder optimization level 2
<code>BuilderFlag::kSTRICT_TYPES</code>	Use for all three flags: 1. <code>kREJECT_EMPTY_ALGORITHMS</code> 2. <code>kDIRECT_IO</code> 3. <code>kPREFER_PRECISION_CONSTRAINTS</code>  Note: When removing enum members (for all enums in this list) we will be changing enumeration in the enum to have sequential numbers.
1. <code>EngineCapability::kDEFAULT</code> 2. <code>EngineCapability::kSAFE_DLA</code> 3. <code>EngineCapability::kSAFE_GPU</code>	1. <code>EngineCapability::kSTANDARD</code> 2. <code>EngineCapability::kDLA_STANDALONE</code> 3. <code>EngineCapability::kSAFETY</code>
<code>IAgorithm::getAlgorithmIOInfo()</code>	<code>IAgorithm::getAlgorithmIOInfoByIndex()</code>
<code>IAgorithmIOInfo::getTensorFormat()</code>	The strides, data type, and vectorization information is sufficient to uniquely identify tensor formats.
1. <code>IBuilder::buildEngineWithConfig()</code> 2. <code>IBuilder::destroy()</code> 3. <code>IBuilder::getMaxBatchSize()</code> 4. <code>IBuilder::setMaxBatchSize()</code>	1. <code>IBuilder::buildSerializedNetwork()</code> 2. <code>delete ObjectName</code> 3. Implicit batch is no longer supported 4. Implicit batch is no longer supported
1. <code>IBuilderConfig::destroy()</code>	1. <code>delete ObjectName</code>

C++ API	Superseded API
2. <code>IBuilderConfig::getMaxWorkspaceSize()</code> 3. <code>IBuilderConfig::getMinTimingIterations()</code> 4. <code>IBuilderConfig::setMaxWorkspaceSize()</code> 5. <code>IBuilderConfig::setMinTimingIterations()</code>	2. <code>IBuilderConfig::getMemoryPoolLimit()</code> with <code>MemoryPoolType::kWORKSPACE</code> 3. <code>IBuilderConfig::getAvgTimingIterations()</code> 4. <code>IBuilderConfig::setMemoryPoolLimit()</code> with <code>MemoryPoolType::kWORKSPACE</code> 5. <code>IBuilderConfig::setAvgTimingIterations()</code>
1. <code>IConvolutionLayer::getDilation()</code> 2. <code>IConvolutionLayer::getKernelSize()</code> 3. <code>IConvolutionLayer::getPadding()</code> 4. <code>IConvolutionLayer::getStride()</code> 5. <code>IConvolutionLayer::setDilation()</code> 6. <code>IConvolutionLayer::setKernelSize()</code> 7. <code>IConvolutionLayer::setPadding()</code> 8. <code>IConvolutionLayer::setStride()</code>	1. <code>IConvolutionLayer::getDilationNd()</code> 2. <code>IConvolutionLayer::getKernelSizeNd()</code> 3. <code>IConvolutionLayer::getPaddingNd()</code> 4. <code>IConvolutionLayer::getStrideNd()</code> 5. <code>IConvolutionLayer::setDilationNd()</code> 6. <code>IConvolutionLayer::setKernelSizeNd()</code> 7. <code>IConvolutionLayer::setPaddingNd()</code> 8. <code>IConvolutionLayer::setStrideNd()</code>
1. <code>ICudaEngine::bindingIsInput()</code> 2. <code>ICudaEngine::destroy()</code> 3. <code>ICudaEngine::getBindingBytesPerComponent()</code> 4. <code>ICudaEngine::getBindingComponentsPerElement()</code> 5. <code>ICudaEngine::getBindingDataType()</code> 6. <code>ICudaEngine::getBindingDimensions()</code> 7. <code>ICudaEngine::getBindingFormat()</code> 8. <code>ICudaEngine::getBindingFormatDesc()</code> 9. <code>ICudaEngine::getBindingIndex()</code> 10. <code>ICudaEngine::getBindingName()</code> 11. <code>ICudaEngine::getBindingVectorizedDim()</code> 12. <code>ICudaEngine::getLocation()</code> 13. <code>ICudaEngine::getMaxBatchSize()</code> 14. <code>ICudaEngine::getNbBindings()</code> 15. <code>ICudaEngine::getProfileDimensions()</code> 16. <code>ICudaEngine::getProfileShapeValues()</code> 17. <code>ICudaEngine::hasImplicitBatchDimension()</code> 18. <code>ICudaEngine::isExecutionBinding()</code> 19. <code>ICudaEngine::isShapeBinding()</code>	1. <code>ICudaEngine::getTensorIOMode()</code> 2. <code>delete ObjectName</code> 3. <code>ICudaEngine::getTensorBytesPerComponent()</code> 4. <code>ICudaEngine::getTensorComponentsPerElement()</code> 5. <code>ICudaEngine::getTensorDataType()</code> 6. <code>ICudaEngine::getTensorShape()</code> 7. <code>ICudaEngine::getTensorFormat()</code> 8. <code>ICudaEngine::getTensorFormatDesc()</code> 9. Name-based methods 10. Name-based methods 11. <code>ICudaEngine::getTensorVectorizedDim()</code> 12. <code>ITensor::getLocation()</code> 13. Implicit batch is no longer supported 14. <code>ICudaEngine::getNbIOTensors()</code> 15. <code>ICudaEngine::getProfileShape()</code> 16. <code>ICudaEngine::getShapeValues()</code> 17. Implicit batch is no longer supported 18. No name-based equivalent replacement 19. <code>ICudaEngine::isShapeInferenceIO()</code>

C++ API	Superseded API
1. <code>IDeconvolutionLayer::getKernelSize()</code> 2. <code>IDeconvolutionLayer::getPadding()</code> 3. <code>IDeconvolutionLayer::getStride()</code> 4. <code>IDeconvolutionLayer::setKernelSize()</code> 5. <code>IDeconvolutionLayer::setPadding()</code> 6. <code>IDeconvolutionLayer::setStride()</code>	1. <code>IDeconvolutionLayer::getKernelSizeNd()</code> 2. <code>IDeconvolutionLayer::getPaddingNd()</code> 3. <code>IDeconvolutionLayer::getStrideNd()</code> 4. <code>IDeconvolutionLayer::setKernelSizeNd()</code> 5. <code>IDeconvolutionLayer::setPaddingNd()</code> 6. <code>IDeconvolutionLayer::setStrideNd()</code>
1. <code>IExecutionContext::destroy()</code> 2. <code>IExecutionContext::enqueue()</code> 3. <code>IExecutionContext::enqueueV2()</code> 4. <code>IExecutionContext::execute()</code> 5. <code>IExecutionContext::getBindingDimensions()</code> 6. <code>IExecutionContext::getShapeBinding()</code> 7. <code>IExecutionContext::getStrides()</code> 8. <code>IExecutionContext::setBindingDimensions()</code> 9. <code>IExecutionContext::setInputShapeBinding()</code> 10. <code>IExecutionContext::setOptimizationProfile()</code>	1. <code>delete ObjectName</code> 2. <code>IExecutionContext::enqueueV3()</code> 3. <code>IExecutionContext::enqueueV3()</code> 4. <code>IExecutionContext::executeV2()</code> 5. <code>IExecutionContext::getTensorShape()</code> 6. <code>IExecutionContext::getTensorAddress()</code> or <code>getOutputTensorAddress()</code> 7. <code>IExecutionContext::getTensorStrides()</code> 8. <code>IExecutionContext::setInputShape()</code> 9. <code>IExecutionContext::setInputTensorAddress()</code> or <code>setTensorAddress()</code> 10. <code>IExecutionContext::setOptimizationProfileAsync()</code>
<code>IFullyConnectedLayer</code>	<code>IMatrixMultiplyLayer</code>
<code>IGpuAllocator::free()</code>	<code>IGpuAllocator::deallocate()</code>
<code>IHostMemory::destroy()</code>	<code>delete ObjectName</code>
1. <code>INetworkDefinition::addConvolution()</code> 2. <code>INetworkDefinition::addDeconvolution()</code> 3. <code>INetworkDefinition::addFullyConnected()</code> 4. <code>INetworkDefinition::addPadding()</code> 5. <code>INetworkDefinition::addPooling()</code> 6. <code>INetworkDefinition::addRNNv2()</code> 7. <code>INetworkDefinition::destroy()</code> 8. <code>INetworkDefinition::hasExplicitPrecision()</code> 9. <code>INetworkDefinition::hasImplicitBatchDimension()</code>	1. <code>INetworkDefinition::addConvolutionNd()</code> 2. <code>INetworkDefinition::addDeconvolutionNd()</code> 3. <code>INetworkDefinition::addMatrixMultiply()</code> 4. <code>INetworkDefinition::addPaddingNd()</code> 5. <code>INetworkDefinition::addPoolingNd()</code> 6. <code>INetworkDefinition::addLoop()</code> 7. <code>delete ObjectName</code> 8. Explicit precision support is removed in 10.0 9. Implicit batch support is removed
<code>IOnnxConfig::destroy()</code>	<code>delete ObjectName</code>
1. <code>IPaddingLayer::getPostPadding()</code>	1. <code>IPaddingLayer::getPostPaddingNd()</code>

C++ API	Superseded API
2. <code>IPaddingLayer::getPrePadding()</code> 3. <code>IPaddingLayer::setPostPadding()</code> 4. <code>IPaddingLayer::setPrePadding()</code>	2. <code>IPaddingLayer::getPrePaddingNd()</code> 3. <code>IPaddingLayer::setPostPaddingNd()</code> 4. <code>IPaddingLayer::setPrePaddingNd()</code>
1. <code>IPoolingLayer::getPadding()</code> 2. <code>IPoolingLayer::getStride()</code> 3. <code>IPoolingLayer::getWindowSize()</code> 4. <code>IPoolingLayer::setPadding()</code> 5. <code>IPoolingLayer::setStride()</code> 6. <code>IPoolingLayer::setWindowSize()</code>	1. <code>IPoolingLayer::getPaddingNd()</code> 2. <code>IPoolingLayer::getStrideNd()</code> 3. <code>IPoolingLayer::getWindowSizeNd()</code> 4. <code>IPoolingLayer::setPaddingNd()</code> 5. <code>IPoolingLayer::setStrideNd()</code> 6. <code>IPoolingLayer::setWindowSizeNd()</code>
<code>IRefitter::destroy()</code>	<code>delete ObjectName</code>
1. <code>IResizeLayer::getAlignCorners()</code> 2. <code>IResizeLayer::setAlignCorners()</code>	1. <code>IResizeLayer::getAlignCornersNd()</code> 2. <code>IResizeLayer::setAlignCornersNd()</code>
1. <code>IRuntime::deserializeCudaEngine(void const* blob, std::size_t size, IPluginFactory* pluginFactory)</code> 2. <code>IRuntime::destroy()</code>	1. Use <code>deserializeCudaEngine</code> with two parameters 2. <code>delete ObjectName</code>
<code>IRNNv2Layer</code>	<code>ILoop</code>
<code>kNV_TENSORRT_VERSION_IMPL</code>	<pre>#define NV_TENSORRT_VERSION_INT(major, minor, patch) ((major) *10000L + (minor) *100L + (patch) *1L)</pre>  Note: TensorRT version encoding was changed to accommodate a two digit minor version.
1. <code>NetworkDefinitionCreationFlag::kEXPLICIT_BATCH</code> 2. <code>NetworkDefinitionCreationFlag::kEXPLICIT_PRECISION</code>	Support is removed in 10.0
1. <code>NV_TENSORRT_SONAME_MAJOR</code> 2. <code>NV_TENSORRT_SONAME_MINOR</code> 3. <code>NV_TENSORRT_SONAME_PATCH</code>	1. <code>NV_TENSORRT_MAJOR</code> 2. <code>NV_TENSORRT_MINOR</code> 3. <code>NV_TENSORRT_PATCH</code>
1. <code>PaddingMode::kCAFFE_ROUND_DOWN</code> 2. <code>PaddingMode::kCAFFE_ROUND_UP</code>	Caffe is not supported since 9.0

C++ API	Superseded API
1. <code>PreviewFeature::kDISABLE_EXTERNAL_TACTICS</code> 2. <code>PreviewFeature::kFASTER_DYNAMIC_SHAPES</code>	1. External tactics are always disabled for core code 2. This flag is on by default
1. <code>ProfilingVerbosity::kDEFAULT</code> 2. <code>ProfilingVerbosity::kVERBOSE</code>	1. <code>ProfilingVerbosity::kLAYER_NAMES_ONLY</code> 2. <code>ProfilingVerbosity::kDETAILED</code>
<code>ResizeMode</code>	Use <code>InterpolationMode</code> , alias is removed
1. <code>RNNDirection</code> 2. <code>RNNGateType</code> 3. <code>RNNInputMode</code> 4. <code>RNNOperation</code>	RNN related data structures are removed
<code>SampleMode::kDEFAULT</code>	<code>SampleMode::kSTRICT_BOUNDS</code>
<code>SliceMode</code>	Use <code>SampleMode</code> , alias is removed

2.5. Removed C++ Plugins

Table 7. Removed C++ Plugins and their Suggested Superseded Plugin

C++ Plugin	Superseded Plugin
1. <code>createAnchorGeneratorPlugin()</code> 2. <code>createBatchedNMSPlugin()</code> 3. <code>createInstanceNormalizationPlugin()</code> 4. <code>createNMSPlugin()</code> 5. <code>createNormalizePlugin()</code> 6. <code>createPriorBoxPlugin()</code> 7. <code>createRegionPlugin()</code> 8. <code>createReorgPlugin()</code> 9. <code>createRPNROIPlugin()</code> 10. <code>createSplitPlugin()</code>	1. <code>GridAnchorPluginCreator::createPlugin()</code> 2. <code>BatchedNMSPluginCreator::createPlugin()</code> 3. <code>InstanceNormalizationPluginCreator::createPlugin()</code> 4. <code>NMSPluginCreator::createPlugin()</code> 5. <code>NormalizePluginCreator::createPlugin()</code> 6. <code>PriorBoxPluginCreator::createPlugin()</code> 7. <code>RegionPluginCreator::createPlugin()</code> 8. <code>ReorgPluginCreator::createPlugin()</code> 9. <code>RPROIPluginCreator::createPlugin()</code> 10. <code>INetworkDefinition::addSlice()</code>
<code>struct Quadruple</code>	Related plugins are removed

2.6. Removed Safety C++ APIs

Table 8. Removed Safety C++ APIs and their Suggested Superseded Safety API

Safety C++ API	Superseded Safety API
1. <code>safe::ICudaEngine::bindingIsInput()</code>	1. <code>safe::ICudaEngine::tensorIOMode()</code>
2. <code>safe::ICudaEngine::getBindingBytesPerComponent()</code>	2. <code>safe::ICudaEngine::getTensorBytesPerComponent()</code>
3. <code>safe::ICudaEngine::getBindingComponentsPerElement()</code>	3. <code>safe::ICudaEngine::getTensorComponentsPerElement()</code>
4. <code>safe::ICudaEngine::getBindingDataType()</code>	4. <code>safe::ICudaEngine::getTensorDataType()</code>
5. <code>safe::ICudaEngine::getBindingDimensions()</code>	5. <code>safe::ICudaEngine::getTensorShape()</code>
6. <code>safe::ICudaEngine::getBindingIndex()</code>	6. <code>safe::name-based methods</code>
7. <code>safe::ICudaEngine::getBindingName()</code>	7. <code>safe::name-based methods</code>
8. <code>safe::ICudaEngine::getBindingVectorizedDim()</code>	8. <code>safe::ICudaEngine::getTensorVectorizedDim()</code>
9. <code>safe::ICudaEngine::getNbBindings()</code>	9. <code>safe::ICudaEngine::getNbIOTensors()</code>
10. <code>safe::ICudaEngine::getTensorFormat()</code>	10. <code>safe::ICudaEngine::getBindingFormat()</code>
1. <code>safe::IExecutionContext::enqueueV2()</code>	1. <code>safe::IExecutionContext::enqueueV3()</code>
2. <code>safe::IExecutionContext::getStrides()</code>	2. <code>safe::IExecutionContext::getTensorStrides()</code>

Chapter 3. trtexec

3.1. trtexec Flag Changes

Table 9. Changes to flag workspace and minTiming

TensorRT 8.x	TensorRT 10.0
<pre>trtexec \ --onnx=/path/to/model.onnx \ --saveEngine=/path/to/engine.trt \ --optShapes=input:\$INPUT_SHAPE \ --avgTiming=1 \ --workspace=1024 \ --minTiming=1</pre>	<pre>trtexec \ --onnx=/path/to/model.onnx \ --saveEngine=/path/to/engine.trt \ --optShapes=input:\$INPUT_SHAPE \ --avgTiming=1 \ --memPoolSize=workspace:1024</pre>

3.2. Removed trtexec Flags

Table 10. Removed trtexec Flags and their Suggested Superseded Flag

trtexec Flag	Superseded Flag
<code>--minTiming</code>	<code>avgTiming</code>
<code>--preview=features options:</code> ▶ <code>disableExternalTacticSourcesForCore0805</code> ▶ <code>fasterDynamicShapes0805</code>	N/A
<code>--workspace=N</code>	<code>--memPoolSize=poolspec</code>

3.3. Deprecated trtexec Flags

Table 11. Deprecated trtexec Flags

trtexec Flag
--buildOnly
--explicitPrecision
--heuristic
--nvtxMode

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.

Arm

Arm, AMBA and Arm Powered are registered trademarks of Arm Limited. Cortex, MPCore and Mali are trademarks of Arm Limited. "Arm" is used to represent Arm Holdings plc; its operating company Arm Limited; and the regional subsidiaries Arm Inc.; Arm KK; Arm Korea Limited.; Arm Taiwan Limited; Arm France SAS; Arm Consulting (Shanghai) Co. Ltd.; Arm Germany GmbH; Arm Embedded Technologies Pvt. Ltd.; Arm Norway, AS and Arm Sweden AB.

HDMI

HDMI, the HDMI logo, and High-Definition Multimedia Interface are trademarks or registered trademarks of HDMI Licensing LLC.

BlackBerry/QNX

Copyright © 2020 BlackBerry Limited. All rights reserved.

Trademarks, including but not limited to BLACKBERRY, EMBLEM Design, QNX, AVIAGE, MOMENTICS, NEUTRINO and QNX CAR are the trademarks or registered trademarks of BlackBerry Limited, used under license, and the exclusive rights to such trademarks are expressly reserved.

Google

Android, Android TV, Google Play and the Google Play logo are trademarks of Google, Inc.



Trademarks

NVIDIA, the NVIDIA logo, and BlueField, CUDA, DALI, DRIVE, Hopper, JetPack, Jetson AGX Xavier, Jetson Nano, Maxwell, NGC, Nsight, Orin, Pascal, Quadro, Tegra, TensorRT, Triton, Turing and Volta are trademarks and/or registered trademarks of NVIDIA Corporation in the United States and other countries. Other company and product names may be trademarks of the respective companies with which they are associated.

Copyright

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

