



NVIDIA TensorRT Migration Guide

Release 6.x and 7.x for NVIDIA DriveOS

NVIDIA Corporation

Jun 25, 2026

Contents

1	Migration Guide Revision History	2
1.1	How to Use This Guide	2
2	Migrating from TensorRT 10.x to 11.x	4
2.1	Strongly Typed Networks Are Now Required	4
2.2	V3 Is the Recommended Plugin Interface	5
2.3	V2 API Methods Replaced with Updated Versions	6
2.4	Other Removals	6
3	Migrating Python Code from TensorRT 10.x to 11.x	8
3.1	Migrating from Weak Typing to Strong Typing	8
3.1.1	Before (TensorRT 10.x)	8
3.1.2	After (TensorRT 11.x)	9
3.1.3	Summary of Changes	9
3.2	Migrating INT8 Calibration to Explicit Quantization	10
3.2.1	Before (TensorRT 10.x)	10
3.2.2	After (TensorRT 11.x)	11
3.2.3	Summary of Changes	11
3.3	Migrating Plugins from V2 to V3	11
3.3.1	Before (TensorRT 10.x)	12
3.3.2	After (TensorRT 11.x)	13
3.3.3	Summary of Changes	14
3.3.4	Known Issues When Migrating Plugins	15
3.4	Migrating Weight Streaming APIs	15
3.4.1	Before (TensorRT 10.x)	15
3.4.2	After (TensorRT 11.x)	16
3.4.3	Summary of Changes	16
3.5	Migrating Memory Management APIs	16
3.5.1	Before (TensorRT 10.x)	16
3.5.2	After (TensorRT 11.x)	16
3.5.3	Summary of Changes	16
3.6	Removed Python APIs and Replacements	17
4	Migrating C++ Code from TensorRT 10.x to 11.x	20
4.1	Migrating from Weak Typing to Strong Typing	20
4.1.1	Before (TensorRT 10.x)	20
4.1.2	After (TensorRT 11.x)	21
4.1.3	Summary of Changes	21
4.2	Migrating INT8 Calibration to Explicit Quantization	21
4.2.1	Before (TensorRT 10.x)	22
4.2.2	After (TensorRT 11.x)	22
4.2.3	Summary of Changes	23
4.3	Migrating Plugins from IPluginV2DynamicExt to IPluginV3	23

4.3.1	Before (TensorRT 10.x - IPluginV2DynamicExt)	24
4.3.2	After (TensorRT 11.x - IPluginV3)	26
4.3.3	Summary of Changes	29
4.3.4	Known Issues When Migrating Plugins	29
4.4	Migrating Weight Streaming APIs	30
4.4.1	Before (TensorRT 10.x)	30
4.4.2	After (TensorRT 11.x)	30
4.4.3	Summary of Changes	30
4.5	Migrating Memory Management APIs	31
4.5.1	Before (TensorRT 10.x)	31
4.5.2	After (TensorRT 11.x)	31
4.5.3	Summary of Changes	31
4.6	Removed C++ APIs and Replacements	32
4.7	Removed C++ Plugins and Replacements	36
4.8	Deprecated BERT Plugins	37
5	Migrating trtexec Usage from TensorRT 10.x to 11.x	38
5.1	Migrating Precision Flags	38
5.1.1	Before (TensorRT 10.x)	38
5.1.2	After (TensorRT 11.x)	38
5.1.3	Summary of Changes	39
5.2	Migrating INT8 Calibration Workflow	39
5.2.1	Before (TensorRT 10.x)	39
5.2.2	After (TensorRT 11.x)	39
5.2.3	Summary of Changes	39
5.3	Removed trtexec Flags and Replacements	39
5.4	Deprecated trtexec Flags and Replacements	40
6	Migrating Safety Runtime Code from TensorRT 10.x to 11.x	42
6.1	Adapting to Build-Time Safety Scope Validation	42
6.1.1	BuilderFlag::kSAFETY_SCOPE is Deprecated	42
6.1.1.1	Before (TensorRT 10.x)	43
6.1.1.2	After (TensorRT 11.x)	43
6.1.1.3	Summary of Changes	43
6.1.2	isNetworkSupported() No Longer Validates Safety Scope	43
6.1.2.1	Before (TensorRT 10.x)	43
6.1.2.2	After (TensorRT 11.x)	43
6.1.2.3	Summary of Changes	44
6.1.3	Interpreting Build-Time Safety Build Failures	44
6.1.4	Removed trtexec Flag --restricted	44
6.2	Kernel Checker Tool Improvements	45
6.3	Recommendation to Use Safety Proxy for Early Bring-Up	45
6.4	Deprecated trtexec_safe Flags and Replacements	45
6.5	Adapting to Auxiliary CUDA Stream Support	45
6.5.1	Preserving 10.x Single-Stream Behavior	46
6.5.2	Adopting Auxiliary Streams	46
6.5.3	Using Debug Overlay Filesystem for Safe Engine Building on QNX Safety	47
7	Migrating IEngineInspector Usage from TensorRT 10.x to 11.x	49
7.1	Bindings Field Replaced by I/O Tensors	49
7.1.1	Before (TensorRT 10.x)	49
7.1.2	After (TensorRT 11.x)	49
7.2	Format / DataType Field Split	50
7.2.1	TensorRT 10.x Output (Before)	50

7.2.2	TensorRT 11.x Output (After)	50
8	Migrating TensorRT from 10.x to 11.x on NVIDIA DriveOS	52
8.1	Platform Eligibility	52
8.2	Who Should Migrate	52
8.3	Migration Checklist	53
8.4	Version Compatibility on DriveOS	53
8.5	CUDA Dependency Changes	54
8.6	DLA Considerations	54
8.7	Recommendation for Production Deployments	54
9	Migrating TensorRT from 10.x to 11.x on Jetson/JetPack	55
9.1	Platform Eligibility	55
9.2	Who Should Migrate	55
9.3	Migration Checklist	56
9.4	Version Compatibility on Jetson	56
9.5	CUDA and JetPack Dependency Changes	56
9.6	DLA Considerations	56
9.7	Cross-Compilation	57
9.8	Recommendation for Edge Deployments	57
10	Appendix: Migrating from TensorRT 8.x to 10.x	58
10.1	How to Use This Appendix	58
10.1.1	Suggested Reading Order	58
10.1.2	Next Steps	59
11	Appendix: Migrating Python Code from TensorRT 8.x to 10.x	60
11.1	Migrating I/O Buffer Allocation to Named Tensors	60
11.1.1	Summary of Changes	62
11.2	Migrating from enqueueV2 to enqueueV3 (Python)	62
11.2.1	Summary of Changes	63
11.3	Migrating Engine Builds to build_serialized_network	64
11.3.1	Summary of Changes	64
11.4	Python APIs Added in 10.x	64
11.4.1	Types	64
11.4.2	Methods and Properties	65
11.5	Removed Python APIs and Replacements	65
12	Appendix: Migrating C++ Code from TensorRT 8.x to 10.x	68
12.1	Migrating from enqueueV2 to enqueueV3 (C++)	68
12.1.1	Summary of Changes	70
12.2	Migrating Dims and IShapeLayer to 64-Bit Types	71
12.3	C++ APIs Added in 10.x	71
12.3.1	Enums	71
12.3.2	Types	71
12.3.3	Methods and Properties	72
12.4	Removed C++ APIs and Replacements	73
12.5	Removed C++ Plugins and Replacements	77
13	Appendix: Migrating trtexec Usage from TensorRT 8.x to 10.x	78
13.1	Migrating --workspace and --minTiming Options	78
13.1.1	Summary of Changes	79
13.2	Removed trtexec Flags and Replacements	79
13.3	Deprecated trtexec Flags and Replacements	80

14 Appendix: Migrating Safety Runtime Code from TensorRT 8.x to 10.x	81
14.1 Migrating from TensorRT 8.x Safety Runtime to 10.x	81
14.2 Removed Safety C++ APIs and Replacements	85

This guide covers the API and tooling changes required to migrate NVIDIA TensorRT workloads on NVIDIA DriveOS 6.x and 7.x. It addresses two migration paths: TensorRT 10.x to 11.x (primary) and TensorRT 8.x to 10.x (appendix).

If you are unfamiliar with these changes, refer to the [TensorRT sample code](#) for clarification.

Chapter 1. Migration Guide Revision History

Table 1.1: Document Revision History

Date	Summary of Change
March 31, 2026	Initial draft for NVIDIA TensorRT 11.0.1 for NVIDIA DriveOS 7.2.5
April 1, 2026	Start of review
June 21, 2026	End of review
June 22, 2026	Approval review

Table 1.2: Document Changes

Date	Summary of Change
April 10, 2026	Reorganized the TensorRT 8.x to 10.x migration material under the appendix (with accompanying content updates) to continue support for those still on TensorRT 8.x. Added the new TensorRT 10.x to 11.x migration guide, which includes the C++ API, Python API, trtexec, and Safety Runtime migration chapters.
April 20, 2026	Added additional chapters: <i>Migrating IEngineInspector Usage from TensorRT 10.x to 11.x</i> , <i>Migrating TensorRT from 10.x to 11.x on NVIDIA DriveOS</i> , and <i>Migrating TensorRT from 10.x to 11.x on Jetson/JetPack</i> . Updated the <i>trtexec migration page</i> with additional deprecated flags (<code>--useCudaGraph</code> , <code>--useSpinWait</code> , <code>--noDataTransfers</code> , <code>--separateProfileRun</code>) and their replacements.
May 11, 2026	Removed a stale “NVIDIA TensorRT 11.x is currently in development” attention block from the TensorRT 10.x to 11.x overview, C++ API , trtexec , Python API , and Safety Runtime migration chapters.

1.1. How to Use This Guide

Choose the starting point that matches your current TensorRT version:

Migrating from TensorRT 10.x to 11.x

Start here if you are already on TensorRT 10.x and need to prepare for TensorRT 11.x.

Appendix: Migrating from TensorRT 8.x to 10.x

Start here if you are still on TensorRT 8.x. After completing the 8.x-to-10.x migration, continue to the 10.x-to-11.x section.

Chapter 2. Migrating from TensorRT 10.x to 11.x

TensorRT 11.x introduces several breaking changes that affect how you build engines, run inference, and use plugins. The sections below summarize the key changes. For language-specific and tool-specific migration details, select the page that matches your integration:

Python Integration

Migrate Python inference and engine-build code, and review Python APIs removed in 11.x with their replacements.

C++ Integration

Migrate C++ inference and engine-build code, and review C++ APIs and plugins removed in 11.x with their replacements.

trtexec Command-Line Usage

Migrate `trtexec` command lines, and review removed or deprecated flag replacements.

Safety Runtime

Migrate safety runtime code paths, including build-time safety scope validation and kernel checker changes.

NVIDIA DriveOS

Platform-specific migration guidance for DRIVE Thor and DRIVE Orin, including version compatibility and DLA considerations.

Jetson / JetPack

Platform-specific migration guidance for Jetson Orin and Jetson Thor, including JetPack dependencies and cross-compilation.

IEngineInspector JSON Output

Migrate code that parses `IEngineInspector` JSON output, including the replacement of `Bindings` with `I/O Tensors` and the split of the `Format/DataType` field.

2.1. Strongly Typed Networks Are Now Required

TensorRT 11.x removes all *weak typing* APIs. In TensorRT 10.x, you could enable reduced-precision types on the model level by setting builder flags such as `BuilderFlag::kFP16` or `BuilderFlag::kINT8`, and TensorRT would automatically consider lower-precision kernels. While this made it simple to get higher throughput, it could compromise accuracy and did not provide adequate control to model authors.

Similarly, implicit quantization is not supported in TensorRT 11.x. To quantize a model, you must add quantize and dequantize layers to the network.

The most straightforward way to get results similar to using the `trtexec --fp16` flag or `BuilderFlag::kFP16` is to use ModelOpt's AutoCast feature on your model. Refer to the [ModelOpt AutoCast](#) documentation for more information; however, for a basic conversion from FP32 layers to FP16, run:

```
python -m modelopt.onnx.autocast --onnx_path model.onnx --output_path model_
→fp16.onnx
```

For quantizing a network, run:

```
python -m modelopt.onnx.quantization --onnx_path model.onnx --calibration_data
→data.npz
```

For more control or other advanced quantization topics (such as quantization-aware training), refer to the [ModelOpt Quantization documentation](#).

You can also manually control a layer's I/O types by using `INetworkDefinition::addCast` or adding `addCast` nodes to the ONNX graph.

Similarly, you can apply quantization to a layer in your network using the `INetworkDefinition::addQuantize` and `INetworkDefinition::addDequantize` APIs, or by adding `QuantizeLinear` and `DequantizeLinear` nodes to the ONNX graph manually. This approach works, but determining the layers that must run at high precision and calibrating the quantization scales for optimal accuracy can be challenging. For most models, use ModelOpt to apply quantization effectively.

2.2. V3 Is the Recommended Plugin Interface

The V2 plugin interfaces (`IPluginV2`, `IPluginV2Ext`, `IPluginV2IOExt`, `IPluginV2DynamicExt`, `IPluginCreator`, `IPluginV2Layer`, and `INetworkDefinition::addPluginV2()`) remain present in TensorRT 11.0 but are deprecated; new plugins should use `IPluginV3` with `IPluginCreatorV3One`, and existing V2 plugins should be migrated. TensorRT 11.x removes the following deprecated methods on the V2 interfaces:

- ▶ `IPluginV2Ext::isOutputBroadcastAcrossBatch` and `IPluginV2Ext::canBroadcastInputAcrossBatch` (also inherited by `IPluginV2DynamicExt`) - implicit-batch broadcast helpers; no replacement is needed because implicit batch is unsupported since 10.0.
- ▶ The legacy `IPluginRegistry` overloads that took `IPluginCreator::registerCreator(IPluginCreator&, ...)`, `deregisterCreator(IPluginCreator const&)`, `getPluginCreator(...)`, and `getPluginCreatorList(...)`. Use the `IPluginCreatorInterface`-typed overloads instead.

Key differences between V2 and V3:

- ▶ **Capability-based design:** `IPluginV3` uses separate capability interfaces (`IPluginV3OneCore`, `IPluginV3OneBuild`, `IPluginV3OneRuntime`) instead of a single monolithic class.
- ▶ **Data-dependent output shapes:** V3 supports declaring size tensors using `IEExprBuilder::declareSizeTensor()`, enabling plugins with data-dependent output shapes.
- ▶ **Shape inputs:** `addPluginV3()` accepts both data inputs and shape inputs, unlike `addPluginV2()`, which only accepts data inputs.
- ▶ **Phase-aware creation:** `IPluginCreatorV3One::createPlugin()` receives a TensorRT-Phase parameter, allowing different behavior during build vs. runtime.

2.3. V2 API Methods Replaced with Updated Versions

Several API methods have been replaced with updated versions that use `int64_t` instead of `size_t`, accept additional parameters, or support asynchronous operation:

Category	Removed	Replacement
Device memory size	<code>getDeviceMemorySize()</code>	<code>getDeviceMemorySizeV2()</code> (returns <code>int64_t</code>)
Device memory (profile)	<code>getDeviceMemorySizeForProfile()</code>	<code>getDeviceMemorySizeForProfileV2()</code> (returns <code>int64_t</code>)
Weight streaming budget	<code>setWeightStreamingBudget()</code> / <code>getWeightStreamingBudget()</code>	<code>setWeightStreamingBudgetV2()</code> / <code>getWeightStreamingBudgetV2()</code>
Profile tensor values	<code>getProfileTensorValues()</code>	<code>getProfileTensorValuesV2()</code> (returns <code>int64_t</code> values)
Deserialization	<code>deserializeCudaEngine(IStreamReader&)</code>	<code>deserializeCudaEngine(IStreamReaderV2&)</code>

2.4. Other Removals

- **Tactic sources:** The `kCUBLAS`, `kCUBLAS_LT`, and `kCUDNN` values of the `TacticSource` enum have been removed. TensorRT no longer uses these external libraries for core operations.
- **Implicit batch:** `allInputShapesSpecified()` has been removed.
- **Built-in plugins removed:** 16 plugins deprecated before TensorRT 10 have been removed, including the NMS family (`BatchedNMS_TRT`, `BatchedNMSDynamic_TRT`, `EfficientNMS_ONNX_TRT`, `NMS_TRT`, `NMSDynamic_TRT`), activations (`Clip_TRT`, `CustomGeluPluginDynamic`, `LReLU_TRT`), and `BatchTilePlugin_TRT`, `CoordConvAC`, `Normalize_TRT`, `Proposal`, `SingleStepLSTMPugin`, `SpecialSlice_TRT`, `Split`. Refer to [Removed C++ Plugins and Replacements](#) for replacements.
- **BERT plugin family deprecated:** The OSS BERT plugin classes (`bertQKVToContextPlugin` / `CustomQKVToContextPluginDynamic` versions 1 through 4, and `CustomEmbLayerNormPluginDynamic`) are deprecated in 11.0.0 and scheduled for removal in a future release. Use the native attention path (`addAttentionV2`) for QKV-to-context, and standard `IGatherLayer` + `addNormalizationV2()` for embedding + layer normalization. Refer to [Deprecated BERT Plugins](#) for per-plugin replacements.
- **ONNX parser:** `supportsModel()` and `parseWithWeightDescriptors()` have been removed.
- **Layer overloads:** Older overloads of `addTopK`, `addNonZero`, and `addNMS` that did not accept an `indicesType` parameter have been removed. Use the versions that accept a `DataType` parameter for the indices output type.
- **Multi-Device preview flag:** The `PreviewFeature::kMULTIDEVICE_RUNTIME_10_16` value has been removed. `setPreviewFeature()` is no longer required to access [Multi-Device Inference](#) features.

Note

If you are migrating from TensorRT 8.x, refer to the [appendix](#) first, then return here for the 10.x-to-11.x changes.

Chapter 3. Migrating Python Code from TensorRT 10.x to 11.x

This page describes how to update Python code when you migrate from TensorRT 10.x to 11.x: paired examples for strong typing, explicit quantization, plugin migration, and updated runtime APIs, followed by lists of Python APIs added and removed in 11.x.

Note

The Python API is not supported on QNX.

Warning

Engine files are executable artifacts that contain compiled CUDA tactics. The `deserialize_cuda_engine` examples in [Migrating from Weak Typing to Strong Typing](#) assume the engine file was built by you on a trusted host or received over a trusted, authenticated channel. Never deserialize an engine file from an untrusted source. Refer to [Engine Deserialization and the Trust Boundary](#) for the full trust-boundary model.

3.1. Migrating from Weak Typing to Strong Typing

TensorRT 11.x removes all precision-enabling builder flags such as `BuilderFlag.FP16` and `BuilderFlag.INT8`. Use ModelOpt AutoCast to convert your ONNX model to mixed precision before building.

3.1.1. Before (TensorRT 10.x)

```
1 import tensorrt as trt
2
3 logger = trt.Logger(trt.Logger.WARNING)
4 builder = trt.Builder(logger)
5 network = builder.create_network()
6 config = builder.create_builder_config()
```

(continues on next page)

(continued from previous page)

```

7
8 # Weak typing: TensorRT automatically considers FP16 kernels
9 config.set_flag(trt.BuilderFlag.FP16)
10
11 parser = trt.OnnxParser(network, logger)
12 with open("model.onnx", "rb") as f:
13     parser.parse(f.read())
14
15 engine_bytes = builder.build_serialized_network(network, config)

```

In TensorRT 11.x, the BuilderFlag.FP16 flag has been removed along with all other weak typing flags. Use ModelOpt AutoCast to convert your FP32 ONNX model to mixed precision before building, then build with a strongly typed network.

3.1.2. After (TensorRT 11.x)

```

1 import tensorrt as trt
2 import modelopt.onnx.autocast as autocast
3 import onnx
4
5 # Step 1: Convert FP32 model to mixed-precision FP16 using ModelOpt AutoCast
6 converted_model = autocast.convert_to_mixed_precision(
7     onnx_path="model.onnx",
8     low_precision_type="fp16",
9     keep_io_types=True,
10 )
11 onnx.save(converted_model, "model_fp16.onnx")
12
13 # Step 2: Build with strongly typed network (strong typing is always on in 11.x)
14 logger = trt.Logger(trt.Logger.WARNING)
15 builder = trt.Builder(logger)
16 network = builder.create_network()
17 config = builder.create_builder_config()
18
19 parser = trt.OnnxParser(network, logger)
20 with open("model_fp16.onnx", "rb") as f:
21     parser.parse(f.read())
22
23 engine_bytes = builder.build_serialized_network(network, config)

```

3.1.3. Summary of Changes

- ▶ Removed `config.set_flag(trt.BuilderFlag.FP16)` and all other precision-enabling builder flags
- ▶ Added a preprocessing step using ModelOpt AutoCast to convert the model to mixed precision
- ▶ The STRONGLY_TYPED network definition creation flag is no longer needed - all networks are strongly typed by default

3.2. Migrating INT8 Calibration to Explicit Quantization

TensorRT 11.x removes `IInt8Calibrator` and all its subclasses. Use `ModelOpt` or manual Q/DQ nodes for explicit quantization instead of runtime calibration.

3.2.1. Before (TensorRT 10.x)

```

1  import tensorrt as trt
2
3  class MyCalibrator(trt.IInt8EntropyCalibrator2):
4      def __init__(self, data_loader):
5          super().__init__()
6          self.data_loader = data_loader
7          self.batch_iter = iter(data_loader)
8
9      def get_batch_size(self):
10         return 1
11
12     def get_batch(self, names):
13         try:
14             batch = next(self.batch_iter)
15             # Copy batch to device and return device pointers
16             return [int(batch.data_ptr())]
17         except StopIteration:
18             return None
19
20     def read_calibration_cache(self):
21         return None
22
23     def write_calibration_cache(self, cache):
24         pass
25
26     logger = trt.Logger(trt.Logger.WARNING)
27     builder = trt.Builder(logger)
28     network = builder.create_network()
29     config = builder.create_builder_config()
30
31     # Implicit quantization via calibrator
32     config.set_flag(trt.BuilderFlag.INT8)
33     config.int8_calibrator = MyCalibrator(data_loader)
34
35     parser = trt.OnnxParser(network, logger)
36     with open("model.onnx", "rb") as f:
37         parser.parse(f.read())
38
39     engine_bytes = builder.build_serialized_network(network, config)

```

In TensorRT 11.x, `IInt8Calibrator` and all its subclasses have been removed. Use `ModelOpt` or manual Q/DQ nodes for explicit quantization.

3.2.2. After (TensorRT 11.x)

```

1  import tensorrt as trt
2
3  # Step 1: Quantize the model offline using ModelOpt (recommended)
4  # python -m modelopt.onnx.quantization --onnx_path model.onnx --calibration_
   ↪ data data.npz
5
6  # Step 2: Build the quantized model (Q/DQ nodes are already in the ONNX graph)
7  logger = trt.Logger(trt.Logger.WARNING)
8  builder = trt.Builder(logger)
9  network = builder.create_network()
10 config = builder.create_builder_config()
11
12 parser = trt.OnnxParser(network, logger)
13 with open("model_quantized.onnx", "rb") as f:
14     parser.parse(f.read())
15
16 engine_bytes = builder.build_serialized_network(network, config)

```

3.2.3. Summary of Changes

- ▶ Removed the `Int8Calibrator` subclass and all calibration-related code
- ▶ Removed `config.set_flag(trt.BuilderFlag.INT8)` and `config.int8_calibrator` assignment
- ▶ Quantization is now applied to the model itself (using Q/DQ nodes) before building, not during building

3.3. Migrating Plugins from V2 to V3

The V2 plugin interfaces have been deprecated since TensorRT 10 and remain present in TensorRT 11.x; no V2 classes are removed in this release, but several deprecated methods on them have been removed (the implicit-batch broadcast helpers on `IPluginV2Ext` / `IPluginV2DynamicExt` and the legacy `IPluginRegistry` overloads that took `IPluginCreator`). New plugins should use the `IPluginV3` interface with `IPluginCreatorV3One`, and existing V2 plugins should be migrated. Refer to the `samples/python/sample_plugin_v2_to_v3_migration` sample that demonstrates how to migrate V2 plugins to V3.

See also

Side-by-Side V2 ↔ V3 API Mapping (in the *TensorRT Developer Guide, Plugins API Migration chapter*)

Method-by-method mapping table grouped by lifecycle phase (core, build, runtime, serialization, network attachment).

Known Migration Issues (in the *TensorRT Developer Guide, Plugins API Migration chapter*)

Known issues encountered when porting V2 plugins, including the empty `PluginField` initializer crash and the strongly-typed network requirement.

Performance: Resolving V2 → V3 Regressions (in the *TensorRT Developer Guide, Plugins API Migration chapter*)

Checklist for resolving performance regressions after migrating a plugin from IPluginV2DynamicExt to IPluginV3.

3.3.1. Before (TensorRT 10.x)

```

1  import tensorrt as trt
2
3  class MyPluginV2(trt.IPluginV2DynamicExt):
4      def __init__(self):
5          super().__init__()
6          self.plugin_namespace = ""
7          self.plugin_type = "MyPlugin"
8          self.plugin_version = "1"
9          self.num_outputs = 1
10
11     def get_output_dimensions(self, output_index, inputs, exprBuilder):
12         return inputs[0] # Same shape as input
13
14     def configure_plugin(self, inp, out):
15         pass
16
17     def supports_format_combination(self, pos, in_out, num_inputs):
18         return in_out[pos].format == trt.TensorFormat.LINEAR and \
19             in_out[pos].type == trt.DataType.FLOAT
20
21     def enqueue(self, input_desc, output_desc, inputs, outputs, workspace,
22 → stream):
23         # Kernel execution
24         pass
25
26     def clone(self):
27         return MyPluginV2()
28
29     def serialize(self):
30         return b""
31
32 # Creator
33 class MyPluginCreatorV2(trt.IPluginCreator):
34     def __init__(self):
35         super().__init__()
36         self.name = "MyPlugin"
37         self.plugin_namespace = ""
38         self.plugin_version = "1"
39         self.field_names = trt.PluginFieldCollection([])
40
41     def create_plugin(self, name, fc):
42         return MyPluginV2()
43
44     def deserialize_plugin(self, name, data):

```

(continues on next page)

(continued from previous page)

```

44         return MyPluginV2()
45
46     # Usage
47     trt.get_plugin_registry().register_creator(MyPluginCreatorV2())
48     plugin = MyPluginV2()
49     layer = network.add_plugin_v2([input_tensor], plugin)

```

3.3.2. After (TensorRT 11.x)

```

1  import tensorrt as trt
2
3  class MyPluginV3(trt.IPluginV3, trt.IPluginV3OneCore,
4                  trt.IPluginV3OneBuild, trt.IPluginV3OneRuntime):
5      def __init__(self):
6          trt.IPluginV3.__init__(self)
7          trt.IPluginV3OneCore.__init__(self)
8          trt.IPluginV3OneBuild.__init__(self)
9          trt.IPluginV3OneRuntime.__init__(self)
10         self.num_outputs = 1
11         self.plugin_namespace = ""
12         self.plugin_name = "MyPlugin"
13         self.plugin_version = "1"
14
15     def get_capability_interface(self, type):
16         return self
17
18     def get_output_data_types(self, input_types):
19         return [input_types[0]] # Same type as input
20
21     def get_output_shapes(self, inputs, shape_inputs, exprBuilder):
22         # Return a list of DimsExprs, one per output
23         output = trt.DimsExprs(len(inputs[0]))
24         for i in range(len(inputs[0])):
25             output[i] = inputs[0][i]
26         return [output]
27
28     def supports_format_combination(self, pos, in_out, num_inputs):
29         return in_out[pos].format == trt.TensorFormat.LINEAR and \
30             in_out[pos].type == trt.DataType.FLOAT
31
32     def configure_plugin(self, inp, out):
33         pass
34
35     def on_shape_change(self, inp, out):
36         return 0
37
38     def enqueue(self, input_desc, output_desc, inputs, outputs, workspace,
39                 ↪ stream):
40         # Kernel execution
41         pass

```

(continues on next page)

(continued from previous page)

```

41
42     def get_fields_to_serialize(self):
43         return trt.PluginFieldCollection([])
44
45     def get_valid_tactics(self):
46         return []
47
48     def set_tactic(self, tactic):
49         return 0
50
51     def clone(self):
52         cloned_plugin = MyPluginV3()
53         cloned_plugin.__dict__.update(self.__dict__)
54         return cloned_plugin
55
56     def attach_to_context(self, context):
57         return self.clone()
58
59 # Creator
60 class MyPluginCreatorV3(trt.IPluginCreatorV3One):
61     def __init__(self):
62         trt.IPluginCreatorV3One.__init__(self)
63         self.name = "MyPlugin"
64         self.plugin_namespace = ""
65         self.plugin_version = "1"
66         self.field_names = trt.PluginFieldCollection([])
67
68     def create_plugin(self, name, fc, phase):
69         return MyPluginV3()
70
71 # Usage
72 trt.get_plugin_registry().register_creator(MyPluginCreatorV3())
73 plugin = MyPluginV3()
74 layer = network.add_plugin_v3([input_tensor], [], plugin)

```

3.3.3. Summary of Changes

- ▶ Plugin class now inherits from `IPluginV3`, `IPluginV3OneCore`, `IPluginV3OneBuild`, and `IPluginV3OneRuntime` instead of `IPluginV2DynamicExt`
- ▶ Added `get_capability_interface()` to return the appropriate capability for each phase
- ▶ `get_output_dimensions()` replaced by `get_output_shapes()`, which receives both data inputs and shape inputs and returns a list of `DimsExprs`
- ▶ `get_output_data_types()` is a new required method
- ▶ `serialize()` replaced by `get_fields_to_serialize()`, which returns a `PluginFieldCollection`
- ▶ Creator inherits from `IPluginCreatorV3One` instead of `IPluginCreator`; `create_plugin()` receives an additional phase parameter and `deserialize_plugin()` is no longer needed

- `network.add_plugin_v2()` replaced by `network.add_plugin_v3()`, which accepts an additional list of shape inputs

3.3.4. Known Issues When Migrating Plugins

- **Empty PluginField initializers can crash V3 dispatch.** When a plugin advertises a `PluginField` whose data is empty (`b""` or a zero-length array), the V3 creator dispatch path can crash during build or deserialization. Populate every entry with a non-empty sentinel value, even when it is unused at runtime:

```
import numpy as np
import tensorrt as trt

# Bad: empty initializer
fields = trt.PluginFieldCollection([
    trt.PluginField("flag", b"", trt.PluginFieldType.INT32),
])

# Good: non-empty sentinel keeps the dispatch path safe
fields = trt.PluginFieldCollection([
    trt.PluginField("flag", np.array([0], dtype=np.int32), trt.
    ↪PluginFieldType.INT32),
])
```

- **Use strongly-typed networks with IPluginV3.** Mixing `IPluginV3` plugins with weakly-typed networks can hit fusion paths that were not exercised by `IPluginV2DynamicExt` and trigger crashes. In TensorRT 11.0.0 all precision-enabling builder flags (`BuilderFlag.FP16`, `INT8`, `BF16`, `FP8`, `INT4`, `FP4`) have been removed, so any network you build is strongly typed by default; no action required for fresh 11.x builds. Authors back-porting V3 plugins to a 10.x build for evaluation must explicitly opt in with `builder.create_network(int(trt.NetworkDefinitionCreationFlag.STRONGLY_TYPED))`.

3.4. Migrating Weight Streaming APIs

The weight streaming API has been updated in TensorRT 11.x. The `minimum_weight_streaming_budget` property has been removed; compute a budget from `streamable_weights_size` and available device memory instead.

3.4.1. Before (TensorRT 10.x)

```
1 import tensorrt as trt
2 import pycuda.driver as cuda
3
4 engine = runtime.deserialize_cuda_engine(engine_bytes)
5
6 # Old API
7 budget = engine.get_minimum_weight_streaming_budget
8 engine.weight_streaming_budget = budget
9 current = engine.weight_streaming_budget
```

3.4.2. After (TensorRT 11.x)

```
1 import tensorrt as trt
2 import pycuda.driver as cuda
3
4 engine = runtime.deserialize_cuda_engine(engine_bytes)
5
6 # V2 API
7 free, total = cuda.mem_get_info()
8 weights_size = engine.streamable_weights_size
9 budget = min(free // 2, weights_size // 2)
10 engine.weight_streaming_budget_v2 = budget
11 current = engine.weight_streaming_budget_v2
```

3.4.3. Summary of Changes

- ▶ `weight_streaming_budget` replaced by `weight_streaming_budget_v2`
- ▶ `get_minimum_weight_streaming_budget` removed - compute your own budget based on available memory and `streamable_weights_size`

3.5. Migrating Memory Management APIs

TensorRT 11.x replaces `device_memory_size` with `device_memory_size_v2` (which returns `int64`) and removes `create_execution_context_without_device_memory()`.

3.5.1. Before (TensorRT 10.x)

```
1 engine = runtime.deserialize_cuda_engine(engine_bytes)
2
3 # Old API
4 mem_size = engine.device_memory_size
5 context = engine.create_execution_context_without_device_memory()
```

3.5.2. After (TensorRT 11.x)

```
1 engine = runtime.deserialize_cuda_engine(engine_bytes)
2
3 # V2 API (returns int64)
4 mem_size = engine.device_memory_size_v2
5 context = engine.create_execution_context()
```

3.5.3. Summary of Changes

- ▶ `device_memory_size` replaced by `device_memory_size_v2` (returns `int64` instead of `size_t`)

- `create_execution_context_without_device_memory()` removed - use `create_execution_context()` with appropriate runtime configuration

3.6. Removed Python APIs and Replacements

Warning

The APIs listed below have been **removed** in TensorRT 11.x and will cause runtime errors if called. Review each entry for its replacement before upgrading.

The following Python APIs have been removed. Each entry shows the removed API and its replacement or migration path.

BuilderFlag.FP16

Strong typing with ModelOpt AutoCast

BuilderFlag.INT8

Explicit quantization with Q/DQ nodes

BuilderFlag.FP8

Explicit quantization with Q/DQ nodes

BuilderFlag.BF16

Strong typing with ModelOpt AutoCast

BuilderFlag.INT4

Explicit quantization with Q/DQ nodes

BuilderFlag.FP4

Explicit quantization with Q/DQ nodes

BuilderFlag.OBEY_PRECISION_CONSTRAINTS

Strong typing (always enforced)

BuilderFlag.PREFER_PRECISION_CONSTRAINTS

Strong typing (always enforced)

BuilderFlag.DIRECT_IO

Removed (unneeded)

IBuilderConfig.int8_calibrator

Explicit quantization with Q/DQ nodes

IBuilderConfig.set_calibration_profile()

Explicit quantization with Q/DQ nodes

IBuilderConfig.get_calibration_profile()

Explicit quantization with Q/DQ nodes

IBuilderConfig.set_quantization_flag()

Explicit quantization with Q/DQ nodes

IBuilderConfig.get_quantization_flag()

Explicit quantization with Q/DQ nodes

IBuilderConfig.quantization_flags

Explicit quantization with Q/DQ nodes

IBuilderConfig.clear_quantization_flag()
Explicit quantization with Q/DQ nodes

ICudaEngine.device_memory_size
ICudaEngine.device_memory_size_v2

ICudaEngine.device_memory_size_for_profile()
ICudaEngine.device_memory_size_for_profile_v2()

ICudaEngine.has_implicit_batch_dimension
Removed (always False)

ICudaEngine.weight_streaming_budget
ICudaEngine.weight_streaming_budget_v2

ICudaEngine.minimum_weight_streaming_budget
Compute from streamable_weights_size and available memory

ICudaEngine.create_execution_context_without_device_memory()
ICudaEngine.create_execution_context()

ICudaEngine.get_profile_tensor_values()
ICudaEngine.get_profile_tensor_values_v2()

IExecutionContext.all_input_shapes_specified
Removed (always True)

IExecutionContext.device_memory
IExecutionContext.device_memory_v2

IInt8Calibrator (all subclasses)
Explicit quantization with Q/DQ nodes

ILayer.precision
Strong typing (set types on tensors directly)

ILayer.precision_is_set
Removed

ILayer.reset_precision()
Removed

ILayer.set_output_type()
Strong typing (set types on tensors directly)

ILayer.output_type_is_set()
Removed

ILayer.reset_output_type()
Removed

INormalizationLayer.compute_precision
Removed (use strong typing)

INetworkDefinition.add_plugin_v2()
INetworkDefinition.add_plugin_v3()

INetworkDefinition.add_normalization()
INetworkDefinition.add_normalization_v2()

IPluginV2DynamicExt
IPluginV3

IPluginCreator

IPluginCreatorV30ne

IPluginRegistry.register_creator() (old overload)

IPluginRegistry.register_creator() (accepts IPluginCreatorInterface)

IPluginRegistry.plugin_creator_list

IPluginRegistry.all_creators

IPluginRegistry.get_plugin_creator()

IPluginRegistry.get_creator()

IRefitter.set_dynamic_range()

Explicit quantization with Q/DQ nodes

IRefitter.get_dynamic_range_min()

Explicit quantization with Q/DQ nodes

IRefitter.get_dynamic_range_max()

Explicit quantization with Q/DQ nodes

IRefitter.get_tensors_with_dynamic_range()

Explicit quantization with Q/DQ nodes

ITensor.set_type()

Strong typing (type determined by network construction)

ITensor.dynamic_range

Explicit quantization with Q/DQ nodes

ITensor.is_dynamic_range_set

Removed

ITensor.reset_dynamic_range()

Removed

TacticSource.CUBLAS

Removed

TacticSource.CUBLAS_LT

Removed

TacticSource.CUDNN

Removed

Chapter 4. Migrating C++ Code from TensorRT 10.x to 11.x

This page describes how to update C++ code when you migrate from TensorRT 10.x to 11.x: paired examples for strongly typed networks, explicit quantization, plugin migration, and updated runtime APIs, followed by lists of C++ APIs added and removed in 11.x.

Warning

Engine files are executable artifacts that contain compiled CUDA tactics. The `deserializeCudaEngine` examples in [Migrating from Weak Typing to Strong Typing](#) assume the engine file was built by you on a trusted host or received over a trusted, authenticated channel. Never deserialize an engine file from an untrusted source. Refer to [Engine Deserialization and the Trust Boundary](#) for the full trust-boundary model.

4.1. Migrating from Weak Typing to Strong Typing

TensorRT 11.x removes all precision-enabling builder flags such as `BuilderFlag::kFP16` and `BuilderFlag::kINT8`. Use ModelOpt AutoCast to convert your ONNX model to mixed precision before building.

4.1.1. Before (TensorRT 10.x)

```
1  auto builder = SampleUniquePtr<nvinfer1::IBuilder>  
    ↳(nvinfer1::createInferBuilder(logger));  
2  auto network = SampleUniquePtr<nvinfer1::INetworkDefinition>(builder->  
    ↳createNetworkV2(0));  
3  auto config = SampleUniquePtr<nvinfer1::IBuilderConfig>(builder->  
    ↳createBuilderConfig());  
4  
5  // Weak typing: TensorRT automatically considers FP16 kernels  
6  config->setFlag(BuilderFlag::kFP16);  
7  
8  auto parser = SampleUniquePtr<nvonnxparser::IParser>(  
9      nvonnxparser::createParser(*network, logger));
```

(continues on next page)

(continued from previous page)

```

10 parser->parseFromFile("model.onnx", static_cast<int>
    ↳(nvinfer1::ILogger::Severity::kWARNING));
11
12 auto plan = SampleUniquePtr<IHostMemory>(builder->
    ↳buildSerializedNetwork(*network, *config));

```

In TensorRT 11.x, BuilderFlag::kFP16 and all other precision-enabling builder flags have been removed. Use ModelOpt AutoCast to convert the ONNX model to mixed precision before building.

4.1.2. After (TensorRT 11.x)

```

1 // Step 1: Convert model to mixed precision offline using ModelOpt:
2 // python -m modelopt.onnx.autocast --onnx_path model.onnx --output_path model_
    ↳fp16.onnx
3
4 // Step 2: Build with strongly typed network (always on in 11.x)
5 auto builder = SampleUniquePtr<nvinfer1::IBuilder>
    ↳(nvinfer1::createInferBuilder(logger));
6 auto network = SampleUniquePtr<nvinfer1::INetworkDefinition>(builder->
    ↳createNetworkV2(0));
7 auto config = SampleUniquePtr<nvinfer1::IBuilderConfig>(builder->
    ↳createBuilderConfig());
8
9 // No precision flags needed - the model itself specifies types
10
11 auto parser = SampleUniquePtr<nvonnxparser::IParser>(
12     nvonnxparser::createParser(*network, logger));
13 parser->parseFromFile("model_fp16.onnx", static_cast<int>
    ↳(nvinfer1::ILogger::Severity::kWARNING));
14
15 auto plan = SampleUniquePtr<IHostMemory>(builder->
    ↳buildSerializedNetwork(*network, *config));

```

4.1.3. Summary of Changes

- ▶ Removed config->setFlag(BuilderFlag::kFP16) and all other precision flags (kINT8, kFP8, kBF16, kINT4, kFP4)
- ▶ Added an offline preprocessing step using ModelOpt AutoCast to produce a mixed-precision ONNX model
- ▶ No code changes needed for the build path itself beyond removing the flag

4.2. Migrating INT8 Calibration to Explicit Quantization

TensorRT 11.x removes IInt8Calibrator and all its subclasses, along with setInt8Calibrator(). Use ModelOpt or manual Q/DQ nodes for explicit quantization instead.

4.2.1. Before (TensorRT 10.x)

```

1  class MyCalibrator : public nvinfer1::IInt8EntropyCalibrator2
2  {
3  public:
4      int32_t getBatchSize() const noexcept override { return 1; }
5
6      bool getBatch(void* bindings[], char const* names[], int32_t nbBindings)
↳noexcept override
7      {
8          // Fill bindings with calibration data
9          if (mCurrentBatch >= mNumBatches)
10             return false;
11          // ... copy data to GPU
12          mCurrentBatch++;
13          return true;
14      }
15
16      void const* readCalibrationCache(size_t& length) noexcept override {
↳return nullptr; }
17      void writeCalibrationCache(void const* ptr, size_t length) noexcept
↳override {}
18
19  private:
20      int32_t mCurrentBatch{0};
21      int32_t mNumBatches{100};
22  };
23
24  // Usage
25  auto config = SampleUniquePtr<nvinfer1::IBuilderConfig>(builder->
↳createBuilderConfig());
26  config->setFlag(BuilderFlag::kINT8);
27
28  MyCalibrator calibrator;
29  config->setInt8Calibrator(&calibrator);
30
31  auto plan = SampleUniquePtr<IHostMemory>(builder->
↳buildSerializedNetwork(*network, *config));

```

In TensorRT 11.x, IInt8Calibrator and all subclasses have been removed along with setInt8Calibrator(). Use ModelOpt or manual Q/DQ nodes.

4.2.2. After (TensorRT 11.x)

```

1  // Step 1: Quantize the model offline using ModelOpt:
2  //   python -m modelopt.onnx.quantization --onnx_path model.onnx --calibration_
↳data data.npz
3  //
4  // Alternatively, add QuantizeLinear/DequantizeLinear nodes to the ONNX graph
↳manually,
5  // or use the INetworkDefinition::addQuantize() and addDequantize() APIs.
6

```

(continues on next page)

(continued from previous page)

```

7 // Step 2: Build the pre-quantized model
8 auto builder = SampleUniquePtr<nvinfer1::IBuilder>
  ↳(nvinfer1::createInferBuilder(logger));
9 auto network = SampleUniquePtr<nvinfer1::INetworkDefinition>(builder->
  ↳createNetworkV2(0));
10 auto config = SampleUniquePtr<nvinfer1::IBuilderConfig>(builder->
  ↳createBuilderConfig());
11
12 auto parser = SampleUniquePtr<nvonnxparser::IParser>(
13     nvonnxparser::createParser(*network, logger));
14 parser->parseFromFile("model_quantized.onnx",
15     static_cast<int>(nvinfer1::ILogger::Severity::kWARNING));
16
17 auto plan = SampleUniquePtr<IHostMemory>(builder->
  ↳buildSerializedNetwork(*network, *config));

```

4.2.3. Summary of Changes

- ▶ Removed the IInt8Calibrator subclass entirely
- ▶ Removed config->setFlag(BuilderFlag::kINT8) and config->setInt8Calibrator()
- ▶ Quantization is applied to the model offline (Q/DQ nodes in the ONNX graph) or using the addQuantize() / addDequantize() network definition APIs

4.3. Migrating Plugins from IPluginV2DynamicExt to IPluginV3

The following example shows a complete plugin migration from V2 to V3 using a NonZero plugin that computes the indices of non-zero elements. This demonstrates V3's support for data-dependent output shapes, which was not possible with V2.

See also

Side-by-Side V2 ↔ V3 API Mapping (in the *TensorRT Developer Guide, Plugins API Migration chapter*)

Method-by-method mapping table grouped by lifecycle phase (core, build, runtime, serialization, network attachment).

Known Migration Issues (in the *TensorRT Developer Guide, Plugins API Migration chapter*)

Known issues encountered when porting V2 plugins, including the empty PluginField initializer crash and the strongly-typed network requirement.

Performance: Resolving V2 ↔ V3 Regressions (in the *TensorRT Developer Guide, Plugins API Migration chapter*)

Checklist for resolving performance regressions after migrating a plugin from IPluginV2DynamicExt to IPluginV3.

4.3.1. Before (TensorRT 10.x - IPluginV2DynamicExt)

```

1  class NonZeroPluginV2 : public nvinfer1::IPluginV2DynamicExt
2  {
3  public:
4      // IPluginV2 core methods
5      char const* getPluginType() const noexcept override { return
↪ "NonZeroPlugin"; }
6      char const* getPluginVersion() const noexcept override { return "1"; }
7      int32_t getNbOutputs() const noexcept override { return 1; }
8
9      // Output dimensions - limited to expressions of input dimensions only
10     DimsExrs getOutputDimensions(int32_t outputIndex, DimsExrs const*
↪ inputs,
11     int32_t nbInputs, IExprBuilder& exprBuilder) noexcept override
12     {
13         // Cannot express data-dependent shapes - must use an upper bound
14         DimsExrs output;
15         output.nbDims = 2;
16         output.d[0] = exprBuilder.operation(DimensionOperation::kPROD,
17         *inputs[0].d[0], *inputs[0].d[1]); // Upper bound: R * C
18         output.d[1] = exprBuilder.constant(2);
19         return output;
20     }
21
22     bool supportsFormatCombination(int32_t pos, PluginTensorDesc const*
↪ inOut,
23     int32_t nbInputs, int32_t nbOutputs) noexcept override
24     {
25         return inOut[pos].format == TensorFormat::kLINEAR;
26     }
27
28     void configurePlugin(DynamicPluginTensorDesc const* in, int32_t nbInputs,
29     DynamicPluginTensorDesc const* out, int32_t nbOutputs) noexcept
↪ override {}
30
31     int32_t enqueue(PluginTensorDesc const* inputDesc, PluginTensorDesc
↪ const* outputDesc,
32     void const* const* inputs, void* const* outputs,
33     void* workspace, cudaStream_t stream) noexcept override
34     {
35         // Execute kernel
36         return 0;
37     }
38
39     size_t getWorkspaceSize(PluginTensorDesc const* inputs, int32_t nbInputs,
40     PluginTensorDesc const* outputs, int32_t nbOutputs) const noexcept
↪ override
41     {
42         return 0;
43     }
44
45     // Serialization

```

(continues on next page)

(continued from previous page)

```

46     size_t getSerializationSize() const noexcept override { return
→ sizeof(bool); }
47     void serialize(void* buffer) const noexcept override
48     {
49         *reinterpret_cast<bool*>(buffer) = mRowOrder;
50     }
51
52     IPluginV2DynamicExt* clone() const noexcept override
53     {
54         return new NonZeroPluginV2(mRowOrder);
55     }
56
57     // ... other required IPluginV2 methods (destroy, setPluginNamespace, etc.
→ )
58
59 private:
60     bool mRowOrder{true};
61 };
62
63 // V2 Plugin Creator
64 class NonZeroCreatorV2 : public nvinfer1::IPluginCreator
65 {
66 public:
67     char const* getPluginName() const noexcept override { return
→ "NonZeroPlugin"; }
68     char const* getPluginVersion() const noexcept override { return "1"; }
69     PluginFieldCollection const* getFieldNames() noexcept override { return &
→ mFC; }
70
71     IPluginV2* createPlugin(char const* name, PluginFieldCollection const*
→ fc) noexcept override
72     {
73         return new NonZeroPluginV2(/*rowOrder=*/true);
74     }
75
76     IPluginV2* deserializePlugin(char const* name, void const* data,
77         size_t length) noexcept override
78     {
79         bool rowOrder = *reinterpret_cast<bool const*>(data);
80         return new NonZeroPluginV2(rowOrder);
81     }
82
83     // ... other required methods
84
85 private:
86     PluginFieldCollection mFC{};
87 };
88
89 // Usage
90 NonZeroPluginV2 plugin(/*rowOrder=*/true);
91 auto* layer = network->addPluginV2(&inputTensor, 1, plugin);

```

4.3.2. After (TensorRT 11.x - IPluginV3)

```

1  class NonZeroPlugin : public IPluginV3, public IPluginV3OneCore,
2                          public IPluginV3OneBuild, public IPluginV3OneRuntime
3  {
4  public:
5      NonZeroPlugin(bool rowOrder) : mRowOrder(rowOrder) {}
6
7      // IPluginV3 - return the appropriate capability interface
8      IPluginCapability* getCapabilityInterface(PluginCapabilityType type)
9      ↪ noexcept override
10     {
11         if (type == PluginCapabilityType::kBUILD)
12             return static_cast<IPluginV3OneBuild*>(this);
13         if (type == PluginCapabilityType::kRUNTIME)
14             return static_cast<IPluginV3OneRuntime*>(this);
15         return static_cast<IPluginV3OneCore*>(this);
16     }
17
18     // IPluginV3OneCore
19     AsciiChar const* getPluginName() const noexcept override { return
20     ↪ "NonZeroPlugin"; }
21     AsciiChar const* getPluginVersion() const noexcept override { return "1";
22     ↪ }
23     AsciiChar const* getPluginNamespace() const noexcept override { return "
24     ↪ "; }
25
26     // IPluginV3OneBuild
27     int32_t getNbOutputs() const noexcept override { return 2; } // data +
28     ↪ size tensor
29
30     int32_t getOutputDataTypes(DataType* outputTypes, int32_t nbOutputs,
31     DataType const* inputTypes, int32_t nbInputs) const noexcept override
32     {
33         outputTypes[0] = DataType::kINT32; // non-zero indices
34         outputTypes[1] = DataType::kINT64; // size tensor
35         return 0;
36     }
37
38     // Output shapes - V3 supports data-dependent shapes via declareSizeTensor
39     int32_t getOutputShapes(DimsExprs const* inputs, int32_t nbInputs,
40     DimsExprs const* shapeInputs, int32_t nbShapeInputs,
41     DimsExprs* outputs, int32_t nbOutputs,
42     IExprBuilder& exprBuilder) noexcept override
43     {
44         auto upperBound = exprBuilder.operation(DimensionOperation::kPROD,
45         *inputs[0].d[0], *inputs[0].d[1]);
46         auto optValue = exprBuilder.operation(DimensionOperation::kFLOOR_DIV,
47         *upperBound, *exprBuilder.constant(2));
48
49         // Declare a size tensor - enables data-dependent output shapes
50         auto numNonZero = exprBuilder.declareSizeTensor(1, *optValue,
51         ↪ *upperBound);

```

(continues on next page)

(continued from previous page)

```

46     outputs[0].nbDims = 2;
47     outputs[0].d[0] = numNonZero; // Data-dependent dimension
48     outputs[0].d[1] = exprBuilder.constant(2);
49
50     outputs[1].nbDims = 0; // Size tensor is a scalar
51     return 0;
52 }
53
54 bool supportsFormatCombination(int32_t pos, DynamicPluginTensorDesc
55 ↪const* inOut,
56     int32_t nbInputs, int32_t nbOutputs) noexcept override
57 {
58     return inOut[pos].desc.format == TensorFormat::kLINEAR;
59 }
60
61 int32_t configurePlugin(DynamicPluginTensorDesc const* in, int32_t
62 ↪nbInputs,
63     DynamicPluginTensorDesc const* out, int32_t nbOutputs) noexcept
64 ↪override
65 {
66     return 0;
67 }
68
69 size_t getWorkspaceSize(DynamicPluginTensorDesc const* inputs, int32_t
70 ↪nbInputs,
71     DynamicPluginTensorDesc const* outputs, int32_t nbOutputs) const
72 ↪noexcept override
73 {
74     return 0;
75 }
76
77 // IPluginV3OneRuntime
78 int32_t enqueue(PluginTensorDesc const* inputDesc, PluginTensorDesc
79 ↪const* outputDesc,
80     void const* const* inputs, void* const* outputs,
81     void* workspace, cudaStream_t stream) noexcept override
82 {
83     // Execute kernel - same as V2
84     return 0;
85 }
86
87 int32_t onShapeChange(PluginTensorDesc const* in, int32_t nbInputs,
88     PluginTensorDesc const* out, int32_t nbOutputs) noexcept override
89 {
90     return 0;
91 }
92
93 // Serialization - uses PluginFieldCollection instead of raw bytes
94 PluginFieldCollection const* getFieldsToSerialize() noexcept override
95 {
96     mDataToSerialize.clear();
97 }

```

(continues on next page)

(continued from previous page)

```

92         mDataToSerialize.emplace_back("rowOrder", &mRowOrder,
→PluginFieldType::kINT32, 1);
93         mFCToSerialize.nbFields = mDataToSerialize.size();
94         mFCToSerialize.fields = mDataToSerialize.data();
95         return &mFCToSerialize;
96     }
97
98     IPluginV3* attachToContext(IPluginResourceContext* context) noexcept
→override
99     {
100         return clone();
101     }
102
103     IPluginV3* clone() noexcept override
104     {
105         return new NonZeroPlugin(mRowOrder);
106     }
107
108 private:
109     bool mRowOrder{true};
110     std::vector<PluginField> mDataToSerialize;
111     PluginFieldCollection mFCToSerialize{};
112 };
113
114 // V3 Plugin Creator
115 class NonZeroCreator : public nvinfer1::IPluginCreatorV3One
116 {
117 public:
118     NonZeroCreator()
119     {
120         mPluginAttributes.emplace_back("rowOrder", nullptr,
→PluginFieldType::kINT32, 1);
121         mFC.nbFields = mPluginAttributes.size();
122         mFC.fields = mPluginAttributes.data();
123     }
124
125     char const* getPluginName() const noexcept override { return
→"NonZeroPlugin"; }
126     char const* getPluginVersion() const noexcept override { return "1"; }
127     char const* getPluginNamespace() const noexcept override { return ""; }
128     PluginFieldCollection const* getFieldNames() noexcept override { return &
→mFC; }
129
130     // Phase-aware creation - no separate deserializePlugin needed
131     IPluginV3* createPlugin(char const* name, PluginFieldCollection const*
→fc,
132     TensorRTPhase phase) noexcept override
133     {
134         bool rowOrder = true;
135         for (int32_t i = 0; i < fc->nbFields; ++i)
136         {
137             if (std::string_view(fc->fields[i].name) == "rowOrder")

```

(continues on next page)

(continued from previous page)

```

138         rowOrder = *static_cast<bool const*>(fc->fields[i].data);
139     }
140     return new NonZeroPlugin(rowOrder);
141 }
142
143 private:
144     PluginFieldCollection mFC{};
145     std::vector<PluginField> mPluginAttributes;
146 };
147
148 // Usage - addPluginV3 accepts both data inputs and shape inputs
149 NonZeroPlugin plugin(/*rowOrder=*/true);
150 ITensor* inputs[] = {&inputTensor};
151 auto* layer = network->addPluginV3(inputs, 1, nullptr, 0, plugin);

```

4.3.3. Summary of Changes

- ▶ Plugin class inherits from IPluginV3, IPluginV3OneCore, IPluginV3OneBuild, and IPluginV3OneRuntime instead of IPluginV2DynamicExt
- ▶ Added getCapabilityInterface() to return the appropriate interface for each phase (core, build, runtime)
- ▶ getOutputDimensions() replaced by getOutputShapes(), which supports data-dependent output shapes using exprBuilder.declareSizeTensor()
- ▶ Added required getOutputDataTypes() method
- ▶ serialize() / getSerializationSize() replaced by getFieldsToSerialize(), which returns a PluginFieldCollection for structured serialization
- ▶ Added onShapeChange() and attachToContext() methods
- ▶ Creator inherits from IPluginCreatorV3One instead of IPluginCreator; createPlugin() takes a TensorRTPhase parameter, and deserializePlugin() is no longer needed - createPlugin() handles both build and runtime phases
- ▶ addPluginV2(inputs, nbInputs, plugin) replaced by addPluginV3(inputs, nbInputs, shapeInputs, nbShapeInputs, plugin)

4.3.4. Known Issues When Migrating Plugins

- ▶ **Empty PluginField initializers can crash V3 dispatch.** When a plugin advertises a PluginField with a nullptr data pointer and length == 0, the V3 creator dispatch path can dereference the pointer during build or deserialization. Populate every entry with a non-null sentinel buffer, even when the value is unused at runtime:

```

// Bad - empty initializer
mPluginAttributes.emplace_back("flag", nullptr, PluginFieldType::kINT32,
    ↪0);

// Good - non-null sentinel keeps the dispatch path safe
static int32_t kDummy = 0;

```

(continues on next page)

(continued from previous page)

```
mPluginAttributes.emplace_back("flag", &kDummy, PluginFieldType::kINT32,
    ↪1);
```

- **Use strongly-typed networks with IPluginV3.** Mixing IPluginV3 plugins with weakly-typed networks can hit fusion paths that were not exercised by IPluginV2DynamicExt and trigger crashes. In TensorRT 11.0.0 all precision-enabling builder flags (BuilderFlag::kFP16, kINT8, kBF16, kFP8, kINT4, kFP4) have been removed, so any network you build is strongly typed by default; no action required for fresh 11.x builds. Authors backporting V3 plugins to a 10.x build for evaluation must explicitly opt in with createNetworkV2(NetworkDefinitionCreationFlag::kSTRONGLY_TYPED).

4.4. Migrating Weight Streaming APIs

The weight streaming API has been updated in TensorRT 11.x. The `getMinimumWeightStreamingBudget()` method has been removed; compute a budget from `getStreamableWeightsSize()` and available device memory instead.

4.4.1. Before (TensorRT 10.x)

```
1 auto engine = SampleUniquePtr<ICudaEngine>(
2   runtime->deserializeCudaEngine(engineData, engineSize));
3
4 // Old API
5 int64_t minBudget = engine->getMinimumWeightStreamingBudget();
6 engine->setWeightStreamingBudget(minBudget);
7 int64_t currentBudget = engine->getWeightStreamingBudget();
```

4.4.2. After (TensorRT 11.x)

```
1 auto engine = SampleUniquePtr<ICudaEngine>(
2   runtime->deserializeCudaEngine(engineData, engineSize));
3
4 // V2 API
5 size_t freeMem, totalMem;
6 cudaMemGetInfo(&freeMem, &totalMem);
7 int64_t weightsSize = engine->getStreamableWeightsSize();
8 int64_t budget = std::min(static_cast<int64_t>(freeMem / 2), weightsSize /
9 ↪2);
10 engine->setWeightStreamingBudgetV2(budget);
11 int64_t currentBudget = engine->getWeightStreamingBudgetV2();
```

4.4.3. Summary of Changes

- `setWeightStreamingBudget()` replaced by `setWeightStreamingBudgetV2()`
- `getWeightStreamingBudget()` replaced by `getWeightStreamingBudgetV2()`

- `getMinimumWeightStreamingBudget()` removed - compute a budget using `getStreamableWeightsSize()` and available device memory

4.5. Migrating Memory Management APIs

TensorRT 11.x replaces `getDeviceMemorySize()` with `getDeviceMemorySizeV2()` (which returns `int64_t`), removes `createExecutionContextWithoutDeviceMemory()`, and replaces `setDeviceMemory(void*)` with `setDeviceMemoryV2(void*, int64_t)`.

4.5.1. Before (TensorRT 10.x)

```

1  auto engine = SampleUniquePtr<ICudaEngine>(
2      runtime->deserializeCudaEngine(engineData, engineSize));
3
4  // Old APIs
5  size_t memSize = engine->getDeviceMemorySize();
6  auto context = SampleUniquePtr<IExecutionContext>(
7      engine->createExecutionContextWithoutDeviceMemory());
8
9  void* deviceMem;
10 cudaMalloc(&deviceMem, memSize);
11 context->setDeviceMemory(deviceMem);

```

4.5.2. After (TensorRT 11.x)

```

1  auto engine = SampleUniquePtr<ICudaEngine>(
2      runtime->deserializeCudaEngine(engineData, engineSize));
3
4  // V2 APIs - int64_t sizes, explicit size parameter
5  int64_t memSize = engine->getDeviceMemorySizeV2();
6  auto context = SampleUniquePtr<IExecutionContext>(engine->
7      createExecutionContext());
8
9  void* deviceMem;
10 cudaMalloc(&deviceMem, memSize);
11 context->setDeviceMemoryV2(deviceMem, memSize);

```

4.5.3. Summary of Changes

- `getDeviceMemorySize()` replaced by `getDeviceMemorySizeV2()` (returns `int64_t` instead of `size_t`)
- `createExecutionContextWithoutDeviceMemory()` removed - use `createExecutionContext()`
- `setDeviceMemory(void*)` replaced by `setDeviceMemoryV2(void*, int64_t)`, which takes an explicit size parameter

4.6. Removed C++ APIs and Replacements

Warning

The APIs listed below have been **removed** in TensorRT 11.x and will cause compile-time errors if used. Review each entry for its replacement before upgrading.

BuilderFlag::kFP16

Strong typing with ModelOpt AutoCast

BuilderFlag::kINT8

Explicit quantization with Q/DQ nodes

BuilderFlag::kFP8

Explicit quantization with Q/DQ nodes

BuilderFlag::kBF16

Strong typing with ModelOpt AutoCast

BuilderFlag::kINT4

Explicit quantization with Q/DQ nodes

BuilderFlag::kFP4

Explicit quantization with Q/DQ nodes

BuilderFlag::kOBEY_PRECISION_CONSTRAINTS

Strong typing (always enforced)

BuilderFlag::kPREFER_PRECISION_CONSTRAINTS

Strong typing (always enforced)

BuilderFlag::kDIRECT_IO

Removed (not needed in 11.x)

IAlgorithm, IAlgorithmContext, IAlgorithmIOInfo, IAlgorithmSelector,**IAlgorithmVariant**

Use editable mode in ITimingCache instead.

IBuilderConfig::setInt8Calibrator(IInt8Calibrator*)

Explicit quantization with Q/DQ nodes

IBuilder::platformHasFastFp16()

Removed; use strongly typed networks instead of querying platform FP16 support. Third-party code that links against TensorRT 11.x - including the ONNX Runtime TensorRT Execution Provider - must migrate away from this API. ONNX Runtime 1.27+ supports TensorRT 11.x. Refer to the [TensorRT 11.0.0 release notes](#) (Known Issues) and [TensorRT 11.1.0 release notes](#) (Fixed Issues).

IBuilder::platformHasFastInt8()

Removed; use explicit quantization with Q/DQ nodes instead of querying platform INT8 support.

IBuilderConfig::getInt8Calibrator()

Removed

IBuilderConfig::setCalibrationProfile(IOptimizationProfile const*)

Removed

IBuilderConfig::getCalibrationProfile()

Removed

IBuilderConfig::setQuantizationFlags(QuantizationFlags)
Removed

IBuilderConfig::getQuantizationFlags()
Removed

IBuilderConfig::clearQuantizationFlag(QuantizationFlag)
Removed

IBuilderConfig::setQuantizationFlag(QuantizationFlag)
Removed

IBuilderConfig::getQuantizationFlag(QuantizationFlag)
Removed

ICudaEngine::createExecutionContextWithoutDeviceMemory()
ICudaEngine::createExecutionContext()

ICudaEngine::getDeviceMemorySize()
ICudaEngine::getDeviceMemorySizeV2()

ICudaEngine::getDeviceMemorySizeForProfile(int32_t)
ICudaEngine::getDeviceMemorySizeForProfileV2(int32_t)

ICudaEngine::getMinimumWeightStreamingBudget()
Compute from getStreamableWeightsSize()

ICudaEngine::getProfileTensorValues(char const*, int32_t, OptProfileSelector)
ICudaEngine::getProfileTensorValuesV2()

ICudaEngine::getWeightStreamingBudget()
ICudaEngine::getWeightStreamingBudgetV2()

ICudaEngine::hasImplicitBatchDimension()
Removed (always false)

ICudaEngine::setWeightStreamingBudget(int64_t)
ICudaEngine::setWeightStreamingBudgetV2(int64_t)

IExecutionContext::allInputShapesSpecified()
Removed (always true)

IExecutionContext::setDeviceMemory(void*)
IExecutionContext::setDeviceMemoryV2(void*, int64_t)

IGpuAllocator::allocate(uint64_t, uint64_t, AllocatorFlags)
IGpuAllocator::allocateAsync(uint64_t, uint64_t, AllocatorFlags, cudaStream_t)

IGpuAllocator::deallocate(void*)
IGpuAllocator::deallocateAsync(void*, cudaStream_t)

IInt8Calibrator (all subclasses)
Explicit quantization with Q/DQ nodes

ILayer::setPrecision(DataType)
Strong typing (set types on tensors directly)

ILayer::getPrecision()
Removed

ILayer::precisionIsSet()
Removed

ILayer::resetPrecision()
Removed

ILayer::setOutputType(int32_t, DataType)
Strong typing (set types on tensors directly)

ILayer::outputTypeIsSet(int32_t)
Removed

ILayer::resetOutputType(int32_t)
Removed

INetworkDefinition::addAttention(..., bool)
INetworkDefinition::addAttentionV2(..., CausalMaskKind)

INetworkDefinition::addNMS(ITensor&, ITensor&, ITensor&)
INetworkDefinition::addNMS(..., DataType) (4-arg version)

INetworkDefinition::addNonZero(ITensor&)
INetworkDefinition::addNonZero(ITensor&, DataType)

INetworkDefinition::addNormalization(...)
INetworkDefinition::addNormalizationV2(...)

INetworkDefinition::addPluginV2(ITensor* const*, int32_t, IPluginV2&)
INetworkDefinition::addPluginV3(...)

INetworkDefinition::addTopK(ITensor&, TopKOperation, int32_t, uint32_t)
INetworkDefinition::addTopK(..., DataType) (5-arg version)

INormalizationLayer::setComputePrecision(DataType)
Removed (use strong typing)

IOutputAllocator::reallocateOutput(char const*, void*, uint64_t, uint64_t)
IOutputAllocator::reallocateOutputAsync(..., cudaStream_t)

IPluginCreator
IPluginCreatorV30ne

IPluginRegistry::deregisterCreator(IPluginCreator const&)
IPluginRegistry::deregisterCreator(IPluginCreatorInterface const&)

IPluginRegistry::getPluginCreator(...)
IPluginRegistry::getCreator(...)

IPluginRegistry::getPluginCreatorList(int32_t*)
IPluginRegistry::getAllCreators(int32_t*)

IPluginRegistry::registerCreator(IPluginCreator&, ...)
IPluginRegistry::registerCreator(IPluginCreatorInterface&, ...)

IPluginV2DynamicExt
IPluginV3

IPluginV2Ext
IPluginV3

IPluginV2IOExt
IPluginV3

IPluginV2Layer
IPluginV3Layer

IRefitter::setDynamicRange(char const*, float, float)

Explicit quantization with Q/DQ nodes

IRefitter::getDynamicRangeMin(char const*)

Removed

IRefitter::getDynamicRangeMax(char const*)

Removed

IRefitter::getTensorsWithDynamicRange()

Removed

IRuntime::deserializeCudaEngine(IStreamReader&)

IRuntime::deserializeCudaEngine(IStreamReaderV2&)

ITensor::setType(DataType)

Strong typing (type determined by network construction)

ITensor::setDynamicRange(float, float)

Explicit quantization with Q/DQ nodes

ITensor::dynamicRangeIsSet()

Removed

ITensor::resetDynamicRange()

Removed

ITensor::getDynamicRangeMin()

Removed

ITensor::getDynamicRangeMax()

Removed

ITensor::setBroadcastAcrossBatch(bool)

Removed (implicit batch not supported)

ITensor::getBroadcastAcrossBatch()

Removed (implicit batch not supported)

TacticSource::kCUBLAS

Removed

TacticSource::kCUBLAS_LT

Removed

TacticSource::kCUDNN

Removed

DetectionOutputParameters

Removed

NMSParameters

Removed

CodeTypeSSD

Removed

4.7. Removed C++ Plugins and Replacements

Warning

The plugins listed below have been **removed** in TensorRT 11.x. Using them will cause compilation or linker errors. Review each entry for its replacement before upgrading.

BatchedNMS_TRT

Use `INetworkDefinition::addNMS()`

BatchedNMSDynamic_TRT

Use `INetworkDefinition::addNMS()`

BatchTilePlugin_TRT

Implement with standard TensorRT layers

Clip_TRT

Use `INetworkDefinition::addActivation()` with `kCLIP`

CoordConvAC

Implement with standard TensorRT layers (concatenate coordinate channels with `IConcatenationLayer`, then apply convolution)

CustomGeluPluginDynamic

Use `INetworkDefinition::addActivation()` with `kGELU_ERF` or `kGELU_TANH`

EfficientNMS_ONNX_TRT

Use `INetworkDefinition::addNMS()`

LReLU_TRT

Use `INetworkDefinition::addActivation()` with `kLEAKY_RELU`

NMS_TRT

Use `INetworkDefinition::addNMS()`

NMSDynamic_TRT

Use `INetworkDefinition::addNMS()`

Normalize_TRT

Use `INetworkDefinition::addNormalizationV2()`

Proposal

Implement with standard TensorRT layers

SingleStepLSTMPugin

Use `INetworkDefinition::addLoop()` or standard RNN decomposition

SpecialSlice_TRT

Use `INetworkDefinition::addSlice()`

Split

Use `INetworkDefinition::addSlice()`

4.8. Deprecated BERT Plugins

The following OSS BERT plugin classes are deprecated in 11.0.0 and scheduled for removal in a future release. Migrate to the listed replacements before upgrading beyond 11.x.

bertQKVToContextPlugin / CustomQKVToContextPluginDynamic

Refer to [Migrate to IAttention](#) for more information.

Chapter 5. Migrating trtexec Usage from TensorRT 10.x to 11.x

This page describes how to update trtexec usage when migrating from TensorRT 10.x to 11.x: before/after command examples and lists of removed or deprecated options.

5.1. Migrating Precision Flags

TensorRT 11.x removes all precision flags (--fp16, --int8, --bf16, --fp8, --int4, --best) because all networks are strongly typed. Use ModelOpt AutoCast to produce a mixed-precision model before building.

5.1.1. Before (TensorRT 10.x)

```
1 trtexec \  
2   --onnx=model.onnx \  
3   --saveEngine=engine.plan \  
4   --fp16
```

In TensorRT 11.x, the --fp16 flag and all other precision flags have been removed because all networks are strongly typed. Use ModelOpt AutoCast to produce a mixed-precision model, then build it.

5.1.2. After (TensorRT 11.x)

```
1 # Step 1: Convert model to mixed precision offline  
2 python -m modelopt.onnx.autocast --onnx_path model.onnx --output_path model_  
→ fp16.onnx  
3  
4 # Step 2: Build with trtexec (strong typing is always on)  
5 trtexec \  
6   --onnx=model_fp16.onnx \  
7   --saveEngine=engine.plan
```

5.1.3. Summary of Changes

- ▶ Removed `--fp16` (and `--bf16`, `--fp8`, `--int4`, `--best`)
- ▶ Added an offline preprocessing step using ModelOpt AutoCast
- ▶ `--stronglyTyped` is no longer needed - it is always enabled

5.2. Migrating INT8 Calibration Workflow

TensorRT 11.x removes the `--int8` and `--calib` flags. Quantization is now applied to the model offline using ModelOpt before building.

5.2.1. Before (TensorRT 10.x)

```
1 trtexec \
2   --onnx=model.onnx \
3   --saveEngine=engine.plan \
4   --int8 \
5   --calib=calibration_cache.bin
```

5.2.2. After (TensorRT 11.x)

```
1 # Step 1: Quantize the model offline using ModelOpt
2 python -m modelopt.onnx.quantization \
3   --onnx_path model.onnx \
4   --calibration_data data.npz
5
6 # Step 2: Build with trtexec
7 trtexec \
8   --onnx=model_quantized.onnx \
9   --saveEngine=engine.plan
```

5.2.3. Summary of Changes

- ▶ Removed `--int8` and `--calib` flags
- ▶ Quantization is applied to the model offline, not during engine build

5.3. Removed trtexec Flags and Replacements

Warning

The flags listed below have been **removed** in TensorRT 11.x. Using them will cause `trtexec` to exit with an error. Review each entry for its replacement before upgrading.

- fp16**
Use ModelOpt AutoCast to produce a mixed-precision model
- bf16**
Use ModelOpt AutoCast to produce a mixed-precision model
- int8**
Use ModelOpt quantization to produce a quantized model
- fp8**
Use ModelOpt quantization to produce a quantized model
- int4**
Use ModelOpt quantization to produce a quantized model
- best**
Use ModelOpt AutoCast and/or quantization
- precisionConstraints**
Removed (strong typing is always enforced)
- layerPrecisions**
Removed (use strong typing in the model)
- layerOutputTypes**
Removed (use strong typing in the model)
- calib**
Use ModelOpt quantization offline
- calibProfile**
Removed

5.4. Deprecated trtexec Flags and Replacements

The following trtexec flags have been deprecated in 11.x but are still accepted. Each entry shows the deprecated flag and its replacement.

- stronglyTyped**
Always enabled; flag accepted but has no effect
- inputIOFormats (type+format)**
Format-only specification (type portion ignored)
- outputIOFormats (type+format)**
Format-only specification (type portion ignored)
- useCudaGraph**
Enabled by default; flag accepted but has no effect. Use `--noCudaGraph` to disable CUDA graph.
- useSpinWait**
Enabled by default; flag accepted but has no effect. Use `--noSpinWait` to disable spin wait for CUDA synchronizations.
- noDataTransfers**
Data transfers are disabled by default; flag accepted but has no effect. Use `--includeDataTransfers` to include data transfer latencies in the benchmarking result.

--separateProfileRun

Always enabled; flag accepted but has no effect. When --dumpProfile or --exportProfile is set, trtexec runs end-to-end performance benchmarking first and then per-layer performance benchmarking as a separate run.

Chapter 6. Migrating Safety Runtime Code from TensorRT 10.x to 11.x

This page describes how to update safety runtime code when migrating from TensorRT 10.x to 11.x.

The most significant change is the removal of frontend safety scope validation. In 10.x, the builder performed a static pre-build check against a “Minimal Safety Scope” allowlist before engine building began. In 11.x, this check is removed and the build-time compiler is the sole authority for safety scope enforcement. As a result, safety build errors now appear as tactic failures with layer-level traceback rather than pre-build scope rejections.

This shift affects builder flag usage, the behavior of `isNetworkSupported()`, error handling patterns, and `trtexec` command-line options. The kernel checker tool also gains MLIR validation in this release. Each section below pairs 10.x and 11.x C++ snippets where applicable.

Important

This section is only applicable when using the TensorRT 11.x safety runtime, which is only available on NVIDIA DriveOS 7.x.

6.1. Adapting to Build-Time Safety Scope Validation

TensorRT 11.x removes pre-build safety scope validation and relies exclusively on build-time enforcement. In 10.x, setting `BuilderFlag::kSAFETY_SCOPE` triggered a pre-compile check against a static “Minimal Safety Scope” allowlist. In 11.x, this pre-build check is removed and the safety scope is now enforced entirely during engine building.

Update your code in three areas: builder configuration, build failure handling, and `isNetworkSupported()` usage.

6.1.1. `BuilderFlag::kSAFETY_SCOPE` is Deprecated

`BuilderFlag::kSAFETY_SCOPE` is retained for API compatibility but has no effect in TensorRT 11.x. Setting it no longer triggers pre-build validation. You can leave existing calls in place or remove them; either way, the flag is silently ignored.

6.1.1.1 Before (TensorRT 10.x)

```

1 config->setFlag(BuilderFlag::kSAFETY_SCOPE);
2 // Triggers static per-layer supportsSafety() checks before engine building.
3 // Returns an error at network definition time for out-of-scope layers.
4 auto serialized = builder->buildSerializedNetwork(*network, *config);

```

6.1.1.2 After (TensorRT 11.x)

```

1 // kSAFETY_SCOPE is now a no-op. Remove it or leave it - it has no effect.
2 // Safety validation occurs exclusively at build time.
3 auto serialized = builder->buildSerializedNetwork(*network, *config);

```

6.1.1.3 Summary of Changes

- BuilderFlag::kSAFETY_SCOPE is deprecated and ignored.
- Safety build failures now surface as build-time tactic errors (for example, No tactic available) with layer-level traceback, rather than pre-build scope errors.

6.1.2. isNetworkSupported() No Longer Validates Safety Scope

In TensorRT 10.x, `IBuilder::isNetworkSupported()` performed static pre-build safety scope checks when `kSAFETY_SCOPE` was set. In 11.x, the function is simplified: it checks only architectural constraint violations (for example, hybrid DLA/GPU mode conflicts and `kSAFETY_SCOPE` flag combinations). A true return value does not guarantee a successful build.

6.1.2.1 Before (TensorRT 10.x)

```

1 // isNetworkSupported() checked static safety scope rules
2 // and could be used as a fast pre-build safety gate.
3 if (!builder->isNetworkSupported(*network, *config)) {
4     // Network was outside the static Minimal Safety Scope.
5     return false;
6 }
7 auto serialized = builder->buildSerializedNetwork(*network, *config);

```

6.1.2.2 After (TensorRT 11.x)

```

1 // isNetworkSupported() checks only architectural constraints (flag
2 // ↪ compatibility,
3 // hybrid DLA/GPU mode). It does NOT validate safety scope.
4 // Use buildSerializedNetwork() for a definitive safety determination.
5 if (!builder->isNetworkSupported(*network, *config)) {
6     // Network has an invalid configuration (e.g., incompatible builder flags).
7     return false;
8 }
9 auto serialized = builder->buildSerializedNetwork(*network, *config);
10 if (!serialized) {

```

(continues on next page)

(continued from previous page)

```

10 // Build failed; build-time safety checks rejected the network.
11 // Inspect error log for layer-level traceback.
12 return false;
13 }

```

6.1.2.3 Summary of Changes

- ▶ `isNetworkSupported()` no longer performs per-layer safety scope checks.
- ▶ If you relied on `isNetworkSupported()` as a definitive safety gate, switch to `buildSerializedNetwork()` for a complete build-time determination.
- ▶ Tensor volume limit and boolean tensor checks have been removed from `isNetworkSupported()` (these are now validated at build-time).

6.1.3. Interpreting Build-Time Safety Build Failures

Because all safety scope enforcement now occurs during engine building, build failures manifest as build-time errors rather than pre-build scope rejections. When a network is outside the certified scope, the builder reports a `No tactic available` error that traces back to the originating layer.

Review your error handling code to process these build-time messages:

```

1 auto serialized = builder->buildSerializedNetwork(*network, *config);
2 if (!serialized) {
3     // Examine the error log for messages such as:
4     // "Autotuner: no tactics to implement operation: ... layers=[ONNX Layer:
5     <LayerName>]"
6     // Use this layer name to identify the unsupported operation and consult
7     the
8     // TensorRT Safety Delta Document for known limitations relative to
9     standard scope.
10 }

```

6.1.4. Removed `trtexec` Flag `--restricted`

Warning

The `--restricted` flag has been **removed** in TensorRT 11.x. Using it will cause `trtexec` to exit with an error.

The `--restricted` `trtexec` flag, which previously enabled pre-build safety scope validation during engine builds, has been removed. Safety restrictions can no longer be applied during standard engine building. Remove `--restricted` from any build scripts or CI pipelines that reference it.

6.2. Kernel Checker Tool Improvements

Important

This section is only applicable when using the TensorRT 11.x safety runtime, which is only available on NVIDIA DriveOS 7.x.

The TensorRT kernel checker tool adds new checks to validate kernels in MLIR form, in addition to the existing checks that validate kernels in CUDA C++ form. Refer to the NVIDIA Deep Learning Inference SEooC 2.2 Safety Tools Manual V0.1 for more information.

6.3. Recommendation to Use Safety Proxy for Early Bring-Up

For safety workflows, begin developing with the TensorRT safety proxy runtime on x86 for early bring-up rather than with the TensorRT standard runtime. The standard and safety runtime APIs differ significantly, as described in the *Migrating Safety Runtime Code from TensorRT 8.x to 10.x* section.

6.4. Deprecated trtexec_safe Flags and Replacements

The following trtexec_safe flags have been deprecated in 11.x but are still accepted. Each entry shows the deprecated flag and its replacement.

--useCudaGraph

Enabled by default; flag accepted but has no effect. A deprecation warning is issued. Use --noCudaGraph to disable CUDA graph usage.

--separateProfileRun

Always enabled; flag accepted but has no effect. A deprecation warning is issued. This flag will be removed in a future release.

6.5. Adapting to Auxiliary CUDA Stream Support

Important

This section is only applicable when using the TensorRT 11.x safety runtime, which is only available on NVIDIA DriveOS 7.x.

The TensorRT 11.x safety runtime now supports auxiliary CUDA streams, lifting the TensorRT 10.x restriction that forced every safety engine to execute on a single stream. By default, the builder may produce engines that require one or more auxiliary streams; when it does, the application is responsible for allocating, registering, and destroying those streams at runtime.

Warning

Existing TensorRT 10.x safety application code that loads a serialized engine and calls `executeAsync()` without first registering auxiliary streams may fail at runtime if the engine was built with `maxAuxStreams > 0`.

There are two migration paths:

- ▶ **Preserve 10.x single-stream behavior** by setting `maxAuxStreams` to 0 at build time. No runtime code changes are required.
- ▶ **Adopt multi-stream execution** by querying the engine's auxiliary-stream count and registering streams before the first `executeAsync()` call.

Note

The standard (non-safety) runtime documents the equivalent APIs in [Within-Inference Multi-Streaming](#). The key safety-specific differences are:

- ▶ The application must allocate and register auxiliary streams; the safety runtime does not auto-create them.
- ▶ The runtime APIs are on `nvinfer2::safe::ITRTGraph` (not `IEExecutionContext`) and live in `NvInferSafeRuntime.h`.

6.5.1. Preserving 10.x Single-Stream Behavior

When `maxAuxStreams` is omitted at build time, the builder may select a non-zero value using internal heuristics. To guarantee single-stream behavior across releases, set the flag explicitly to 0 at build time:

```
1 auto config = builder->createBuilderConfig();
2 config->setMaxAuxStreams(0); // produces a single-stream engine
3 // ... rest of build configuration ...
```

The same flag is available in `trtexec`:

```
trtexec --onnx=model.onnx --maxAuxStreams=0 ...
```

6.5.2. Adopting Auxiliary Streams

If your workflow allows multi-stream engines, the aux-stream-specific steps to add to your runtime code are:

1. Query the number of auxiliary streams the engine expects (after creating the safety graph from the serialized engine).
2. Allocate that many CUDA streams with the `cudaStreamNonBlocking` flag.
3. Register the streams using `setAuxStreams()` before the first call to `executeAsync()`. This is an INIT-phase operation (before any inference launches).
4. Destroy the streams after the final `sync()` call.

If `getNbAuxStreams()` returns 0, the engine is single-stream and no auxiliary-stream handling is required - you can skip the allocation and registration steps.

The relevant API lives in namespace `nvinfer2::safe` in `NvInferSafeRuntime.h`:

```

1 // 1. Create the safety graph from the serialized engine.
2 nvinfer2::safe::ITRTGraph* graph = nullptr;
3 nvinfer2::safe::createTRTGraph(graph, planData, planSize, recorder, true);
4
5 // 2. Query the number of auxiliary streams this engine requires.
6 int32_t nbAuxStreams = 0;
7 graph->getNbAuxStreams(nbAuxStreams);
8
9 // 3. Allocate that many non-blocking streams. The application owns
10 //    the lifetime of these streams.
11 std::vector<cudaStream_t> auxStreams(nbAuxStreams);
12 for (int32_t i = 0; i < nbAuxStreams; ++i)
13 {
14     cudaStreamCreateWithFlags(&auxStreams[i], cudaStreamNonBlocking);
15 }
16
17 // 4. Register the streams BEFORE the first executeAsync call.
18 //    The count must exactly match getNbAuxStreams().
19 graph->setAuxStreams(auxStreams.data(), nbAuxStreams);
20
21 // 5. Create the main inference stream.
22 cudaStream_t mainStream;
23 cudaStreamCreateWithFlags(&mainStream, cudaStreamNonBlocking);
24
25 // 6. Run inference.
26 graph->executeAsync(mainStream);
27
28 // 7. Wait for the full inference to complete.
29 graph->sync();
30
31 // 8. Cleanup - must occur AFTER the final sync().
32 for (auto s : auxStreams)
33 {
34     cudaStreamDestroy(s);
35 }
36 cudaStreamDestroy(mainStream);
37 nvinfer2::safe::destroyTRTGraph(graph);

```

6.5.3. Using Debug Overlay Filesystem for Safe Engine Building on QNX Safety

Starting in TensorRT 11.0.1, components of the TensorRT safety builder are shipped in the debug overlay filesystem on QNX Safety. Specifically, the `timing_server` executable and certain dependent libraries are mounted in the debug overlay filesystem in the `/nvidia_overlay` directory rather than `/usr/libnvidia`, where they previously shipped. Other platforms, such as QNX Standard and x86, are unchanged.

To mount the debug overlay, specify `debug_sr` as the platform configuration table (PCT) argument to

bind_partitions, as shown in the *TensorRT 11.0.1 Installation Guide for DriveOS 7.2.5*:

```
./make/bind_partitions -b <board_variant> qnx -p debug_sr
```

The debug overlay is not needed for production runtime. The `libnvinfer_safe.so*` libraries are still available at `/usr/libnvidia`.

Similarly, the `--remoteAutoTuningConfig` flag to `trtexec` should refer to `/nvidia_overlay`:

```
--remoteAutoTuningConfig="ssh://root:root@${TARGET_IP}:22?remote_exec_path=/  
→nvidia_overlay/timing_server&remote_lib_path=/nvidia_overlay:/usr/libnvidia&  
→dump_remote_stdout=on&dump_remote_stderr=on&verbose=on"
```

Chapter 7. Migrating IEngineInspector Usage from TensorRT 10.x to 11.x

If your application parses the JSON output of `IEngineInspector::getLayerInformation()` or `IEngineInspector::getEngineInformation()`, the output schema has changed in TensorRT 11.x. Update your parsing logic as described below.

For general Engine Inspector usage and API details, refer to [Engine Inspector](#).

7.1. Bindings Field Replaced by I/O Tensors

The `Bindings` field, which was a flat list of I/O tensor names as strings, has been removed. It is replaced by the new `I/O Tensors` field, which is a list of JSON objects. Each object contains structured information about the tensor, including its name, I/O mode, data type, dimensions, memory location, and whether it is a shape inference I/O.

7.1.1. Before (TensorRT 10.x)

```
{
  "Bindings": [
    "input_0",
    "output_0"
  ]
}
```

7.1.2. After (TensorRT 11.x)

```
{
  "I/O Tensors": [
    {
      "Name": "input_0",
      "IOMode": "Input",
      "DataType": "Float",
      "Dimensions": [1, 3, 224, 224],
```

(continues on next page)

(continued from previous page)

```

    "Location": "Device",
    "IsShapeInferenceIO": false,
    "ProfileInfo": [
      {
        "MinShape": [1, 3, 224, 224],
        "OptShape": [1, 3, 224, 224],
        "MaxShape": [1, 3, 224, 224],
        "Format": "Row major linear FP32 format (kLINEAR)"
      }
    ],
    {
      "Name": "output_0",
      "IOMode": "Output",
      "DataType": "Float",
      "Dimensions": [1, 1000],
      "Location": "Device",
      "IsShapeInferenceIO": false,
      "ProfileInfo": [
        {
          "Format": "Row major linear FP32 format (kLINEAR)"
        }
      ]
    }
  ]
}

```

7.2. Format / DataType Field Split

In layer-level information, the combined Format/DataType field has been removed and replaced with two separate fields: Format and DataType. This allows you to read the data type directly without parsing the format string.

7.2.1. TensorRT 10.x Output (Before)

```

{
  "Name": "Conv_0",
  "Inputs": [
    {
      "Name": "input_0",
      "Format/DataType": "Row major linear FP32"
    }
  ]
}

```

7.2.2. TensorRT 11.x Output (After)

```

{
  "Name": "Conv_0",

```

(continues on next page)

(continued from previous page)

```
"Inputs": [  
  {  
    "Name": "input_0",  
    "Datatype": "Float",  
    "Format": "(Linear) Row major linear format"  
  }  
]
```

Chapter 8. Migrating TensorRT from 10.x to 11.x on NVIDIA DriveOS

This page describes platform-specific migration considerations for TensorRT on NVIDIA DriveOS when upgrading from TensorRT 10.x to 11.x. For API-level migration details, refer to the [C++](#), [Python](#), [trtexec](#), and [Safety Runtime](#) migration pages.

Important

TensorRT 11.x is supported on DriveOS 7.2.5 and later. DriveOS 7.0 customers on DRIVE Orin remain on TensorRT 10.x, which is maintained on a separate branch. If you are running DriveOS 7.0 with TensorRT 10.x, no migration is required.

8.1. Platform Eligibility

Platform	TensorRT Version		Notes
DRIVE Thor (DriveOS 7.2.5+)	11.x		Full upgrade path. Follow the migration steps on this page.
DRIVE Orin (DriveOS 7.0)	10.x	(maintained branch)	Hard-forked TensorRT 10.x branch. No upgrade to 11.x.
DRIVE Orin (DriveOS 7.3)	10.x		Out of scope for TensorRT 11.x in this release.

8.2. Who Should Migrate

You need to take action if your DriveOS application does any of the following:

- ▶ Uses implicit INT8 quantization (`BuilderFlag::kINT8` with `IInt8Calibrator`, or `--int8` with `trtexec`).
- ▶ Relies on weakly typed builder flags (`BuilderFlag::kFP16`, `BuilderFlag::kBF16`, `BuilderFlag::kFP8`, or equivalent flags such as `--fp16` with `trtexec`).

- Uses IPluginV2DynamicExt or IPluginCreator.
- References removed tactic sources (kCUBLAS, kCUBLAS_LT, kCUDNN).

If your application already uses strong typing, explicit quantization (Q/DQ nodes), and IPluginV3, the upgrade to TensorRT 11.x should require minimal code changes beyond recompilation.

8.3. Migration Checklist

1. **Migrate from weak typing to strong typing.** TensorRT 11.x removes all precision-enabling builder flags. Use [ModelOpt AutoCast](#) to convert ONNX models to mixed precision before building. Refer to the [Strongly Typed Networks](#) section for details.
2. **Migrate from implicit INT8 to explicit quantization.** The IInt8Calibrator class and all subclasses have been removed. Use [ModelOpt Quantization](#) to produce models with Q/DQ nodes. Refer to the [C++](#) or [Python](#) migration page for before/after examples.
3. **Migrate plugins from V2 to V3.** IPluginV2DynamicExt and IPluginCreator have been deprecated. All plugins must use IPluginV3 with IPluginCreatorV3One. Refer to the [C++ plugin migration](#) section for a complete NonZero plugin example.
4. **Update V2 API calls to V2 replacements.** Several methods (getDeviceMemorySize, setWeightStreamingBudget, setDeviceMemory) have been replaced with V2 versions that use int64_t or accept additional parameters. Refer to the [V2 API replacements table](#).
5. **Update safety runtime code (if applicable).** Pre-build safety scope validation has been removed. Build failures now surface as tactic errors with layer-level traceback. Refer to the [Safety Runtime migration](#) page.
6. **Remove `--restricted` from `trtexec` scripts.** The `--restricted` flag has been removed in TensorRT 11.x. Remove it from any build scripts or CI pipelines.
7. **Rebuild all engines.** Engines built with TensorRT 10.x using the `BuilderFlag::kVERSION_COMPATIBLE` setting are forward-compatible with the TensorRT 11.x runtime. However, to take advantage of 11.x optimizations and new features, rebuild engines with the 11.x builder.

8.4. Version Compatibility on DriveOS

TensorRT 11.x maintains forward compatibility with TensorRT 10.x engines:

- Engines built with TensorRT 10.x can run on the TensorRT 11.x runtime.
- Engines built with TensorRT 8.x or 9.x are **not** compatible with the TensorRT 11.x runtime. TensorRT 8.x uses CUDA 11.x, which is not supported by TensorRT 11.x.
- Backward compatibility is maintained within a major release. For example, a version-compatible engine built with TensorRT 11.5 runs on TensorRT 11.5 or any later 11.x release, but not on TensorRT 11.0 through 11.4, and not on TensorRT 10.x.

If your DriveOS deployment currently uses TensorRT 8.x, you must first migrate to TensorRT 10.x using the [8.x to 10.x appendix](#), then follow this guide for the 10.x to 11.x migration.

8.5. CUDA Dependency Changes

TensorRT 11.x requires CUDA 13.3 or later. Verify that your DriveOS BSP includes a compatible CUDA toolkit before upgrading. TensorRT 11.x drops support for CUDA 11.x, which means TensorRT engines built against CUDA 11.x cannot run on the TensorRT 11.x runtime.

8.6. DLA Considerations

DLA is not supported in TensorRT 11.0; DLA support will be re-introduced in a later minor version update.

8.7. Recommendation for Production Deployments

Tip

For automotive customers already in production on non-safety platforms prior to the TensorRT 11.x release, NVIDIA recommends remaining on the TensorRT 10.x branch to ensure optimal stability and performance. Migrate to 11.x when your next development cycle begins.

For safety customers who must upgrade to 11.x, the TensorRT team provides mitigation plans for all removed APIs and will work directly with customers to ensure a smooth transition. Contact your NVIDIA representative for assistance.

Chapter 9. Migrating TensorRT from 10.x to 11.x on Jetson/JetPack

This page describes platform-specific migration considerations for TensorRT on NVIDIA Jetson platforms (JetPack) when upgrading from TensorRT 10.x to 11.x. For API-level migration details, refer to the [C++](#), [Python](#), and [trtexec](#) migration pages.

9.1. Platform Eligibility

Platform	TensorRT Version	Notes
Jetson Thor (JetPack 7.x)	11.x	Full upgrade path. Follow the migration steps on this page.
Jetson Orin (JetPack 6.x)	10.x or 11.x	Optional but recommended upgrade path.

9.2. Who Should Migrate

You need to take action if your Jetson application does any of the following:

- ▶ Uses implicit INT8 quantization (`BuilderFlag::kINT8` with `IInt8Calibrator`, or `--int8` with `trtexec`).
- ▶ Relies on weakly typed builder flags (`BuilderFlag::kFP16`, `BuilderFlag::kBF16`, `BuilderFlag::kFP8`, or equivalent flags such as `--fp16` with `trtexec`).
- ▶ Uses `IPluginV2DynamicExt` or `IPluginCreator`.
- ▶ References removed tactic sources (`kCUBLAS`, `kCUBLAS_LT`, `kCUDNN`).

If your application already uses strong typing, explicit quantization (Q/DQ nodes), and `IPluginV3`, the upgrade to TensorRT 11.x should require minimal code changes beyond recompilation.

9.3. Migration Checklist

1. **Migrate from weak typing to strong typing.** TensorRT 11.x removes all precision-enabling builder flags. Use [ModelOpt AutoCast](#) to convert ONNX models to mixed precision before building. Refer to the [Strongly Typed Networks](#) section for details.
2. **Migrate from implicit INT8 to explicit quantization.** The `IInt8Calibrator` class and all sub-classes have been removed. Use [ModelOpt Quantization](#) to produce models with Q/DQ nodes. Refer to the [C++](#) or [Python](#) migration page for before/after examples.
3. **Migrate plugins from V2 to V3.** `IPluginV2DynamicExt` and `IPluginCreator` have been deprecated. All plugins must use `IPluginV3` with `IPluginCreatorV3One`. Refer to the [C++ plugin migration](#) section for a complete NonZero plugin example.
4. **Update V2 API calls to V2 replacements.** Several methods (`getDeviceMemorySize`, `setWeightStreamingBudget`, `setDeviceMemory`) have been replaced with V2 versions that use `int64_t` or accept additional parameters. Refer to the [V2 API replacements table](#).
5. **Rebuild all engines.** Engines built with TensorRT 10.x using the `BuilderFlag::kVERSION_COMPATIBLE` setting are forward-compatible with the TensorRT 11.x runtime. However, to take advantage of 11.x optimizations and new features, rebuild engines with the 11.x builder.

9.4. Version Compatibility on Jetson

TensorRT 11.x maintains forward compatibility with TensorRT 10.x engines:

- ▶ Engines built with TensorRT 10.x can run on the TensorRT 11.x runtime.
- ▶ Engines built with TensorRT 8.x or 9.x are **not** compatible with the TensorRT 11.x runtime. TensorRT 8.x uses CUDA 11.x, which is not supported by TensorRT 11.x.
- ▶ Backward compatibility is maintained within a major release. For example, a version-compatible engine built with TensorRT 11.5 runs on TensorRT 11.5 or any later 11.x release, but not on TensorRT 11.0 through 11.4, and not on TensorRT 10.x.

If your Jetson deployment currently uses TensorRT 8.x, you must first migrate to TensorRT 10.x using the [8.x to 10.x appendix](#), then follow this guide for the 10.x to 11.x migration.

9.5. CUDA and JetPack Dependency Changes

TensorRT 11.x requires CUDA 13.2 Update 1 or later. Verify that your JetPack version includes a compatible CUDA toolkit before upgrading. Refer to the JetPack release notes for supported CUDA versions.

TensorRT 11.x drops support for CUDA 11.x. Engines built against CUDA 11.x cannot run on the TensorRT 11.x runtime.

9.6. DLA Considerations

DLA is not supported in TensorRT 11.0; DLA support will be re-introduced in a later minor version update.

9.7. Cross-Compilation

If you cross-compile TensorRT applications for Jetson on an x86 host, ensure that both the host TensorRT SDK and the target JetPack installations are upgraded to 11.x. Mixing TensorRT 10.x headers on the host with 11.x libraries on the target (or vice versa) will cause compilation or linker errors due to removed APIs.

9.8. Recommendation for Edge Deployments

Tip

For Jetson applications in production, NVIDIA recommends testing the TensorRT 11.x migration in a staging environment before deploying to production devices. Verify model accuracy, inference latency, and memory usage against your TensorRT 10.x baseline.

Customers still using TensorRT 10.x should begin migrating to strong typing, explicit quantization, and IPluginV3 to reduce the scope of changes when upgrading. The TensorRT team provides mitigation plans for all removed APIs.

Chapter 10. Appendix: Migrating from TensorRT 8.x to 10.x

This appendix explains how to migrate from NVIDIA TensorRT 8.x to 10.x.

If you are unfamiliar with these changes, refer to our [sample code](#) for clarification.

10.1. How to Use This Appendix

Select the migration surface that matches what you ship:

Python Integration

Migrate Python inference and engine-build flows, and review Python APIs added and removed in 10.x.

C++ Integration

Migrate C++ inference and engine-build flows, including 64-bit dimension updates, and review C++ APIs and plugins removed with replacements.

trtexec Command-Line Usage

Migrate `trtexec` command lines, and review removed or deprecated flag replacements.

Safety Runtime

Migrate safety runtime code paths, and review removed safety C++ APIs with replacements.

10.1.1. Suggested Reading Order

1. Start with *Python Integration* or *C++ Integration* for application code paths (read both only if you maintain both stacks).
2. Add *trtexec Command-Line Usage* when your workflow depends on CLI flags or scripts.
3. Read *Safety Runtime* when you target the safety runtime instead of the standard runtime.

Attention

TensorRT 10.x safety proxy runtime support on QNX Standard (without QNX Safety support) is available in limited, alpha-quality capacity in this release. Full functional safety runtime support, including QNX Safety, is planned for a future sideband drop on top of NVIDIA DriveOS 7.2.5.

10.1.2. Next Steps

After completing the 8.x-to-10.x migration, continue to [Migrating from TensorRT 10.x to 11.x](#) for the next set of API and tooling changes.

Chapter 11. Appendix: Migrating Python Code from TensorRT 8.x to 10.x

This page describes how to update Python code when you migrate from TensorRT 8.x to 10.x: paired examples for the name-based tensor API (buffers and I/O), enqueueV3 execution, and build_serialized_network for engine build, followed by lists of Python APIs added and removed in 10.x.

Important

The Python API is not supported on QNX. If you are targeting QNX (Standard or Safety), use the C++ API; refer to [Appendix: Migrating C++ Code from TensorRT 8.x to 10.x](#).

Attention

DriveOS 7.2.5 scope: DLA-related rename entries documented later on this page (such as EngineCapability.kSAFE_DLA → EngineCapability.DLA_STANDALONE) are listed for API completeness only. **DLA is not supported in TensorRT 11.0.1 for NVIDIA DriveOS 7.2.5.** For current DriveOS guidance, refer to [Migrating TensorRT from 10.x to 11.x on NVIDIA DriveOS](#).

11.1. Migrating I/O Buffer Allocation to Named Tensors

TensorRT 10.x replaces the binding-based API with a name-based tensor API. Use num_io_tensors and get_tensor_name() to iterate over I/O tensors, and get_tensor_mode() to check whether a tensor is an input or output.

Before (TensorRT 8.x)

```
1 def allocate_buffers(self, engine):  
2     '''  
3     Allocates all buffers required for an engine; that is, host and device inputs
```

(continues on next page)

(continued from previous page)

```

4  ↪and outputs.
5  '''
6  inputs = []
7  outputs = []
8  bindings = []
9  stream = cuda.Stream()
10
11 # binding is the name of input/output
12 for binding in the engine:
13     size = trt.volume(engine.get_binding_shape(binding)) * engine.max_batch_
14     ↪size
15     dtype = trt.nptype(engine.get_binding_dtype(binding))
16
17     # Allocate host and device buffers
18     host_mem = cuda.pagelocked_empty(size, dtype) # page-locked memory buffer
19     ↪(will not be swapped to disk)
20     device_mem = cuda.mem_alloc(host_mem.nbytes)
21
22     # Append the device buffer address to device bindings.
23     # When cast to int, it is a linear index into the context's memory (like
24     ↪memory address).
25     bindings.append(int(device_mem))
26
27     # Append to the appropriate input/output list.
28     if engine.binding_is_input(binding):
29         inputs.append(self.HostDeviceMem(host_mem, device_mem))
30     else:
31         outputs.append(self.HostDeviceMem(host_mem, device_mem))
32
33 return inputs, outputs, bindings, stream

```

After (TensorRT 10.x)

```

1  def allocate_buffers(self, engine):
2  '''
3  Allocates all buffers required for an engine; that is, host and device inputs
4  ↪and outputs.
5  '''
6  inputs = []
7  outputs = []
8  bindings = []
9  stream = cuda.Stream()
10
11 for i in range(engine.num_io_tensors):
12     tensor_name = engine.get_tensor_name(i)
13     size = trt.volume(engine.get_tensor_shape(tensor_name))
14     dtype = trt.nptype(engine.get_tensor_dtype(tensor_name))
15
16     # Allocate host and device buffers
17     host_mem = cuda.pagelocked_empty(size, dtype) # page-locked memory buffer
18     ↪(will not be swapped to disk)

```

(continues on next page)

(continued from previous page)

```

17     device_mem = cuda.mem_alloc(host_mem.nbytes)
18
19     # Append the device buffer address to device bindings.
20     # When cast to int, it is a linear index into the context's memory (like
    ↪ memory address).
21     bindings.append(int(device_mem))
22
23     # Append to the appropriate input/output list.
24     if engine.get_tensor_mode(tensor_name) == trt.TensorIOMode.INPUT:
25         inputs.append(self.HostDeviceMem(host_mem, device_mem))
26     else:
27         outputs.append(self.HostDeviceMem(host_mem, device_mem))
28
29     return inputs, outputs, bindings, stream

```

11.1.1. Summary of Changes

- Changed from binding-based iteration to name-based API using `num_io_tensors` and `get_tensor_name()`
- Replaced `binding_is_input()` with `get_tensor_mode()` to check I/O mode

Note

The `HostDeviceMem` helper class used in the examples above is a simple container that pairs host (CPU) and device (GPU) memory allocations. It is not part of the TensorRT API. A minimal implementation is:

```
from collections import namedtuple
```

```
HostDeviceMem = namedtuple("HostDeviceMem", ["host", "device"])
```

11.2. Migrating from enqueueV2 to enqueueV3 (Python)

The examples below show TensorRT 8.x first, then TensorRT 10.x, for the same inference task. In TensorRT 10.x, `enqueueV3` replaces `enqueueV2`: call `set_tensor_address` for each I/O tensor before `execute_async_v3`, as shown in the After tab.

Before (TensorRT 8.x)

```

1  # Allocate device memory for inputs.
2  d_inputs = [cuda.mem_alloc(input_nbytes) for binding in range(input_num)]
3
4  # Allocate device memory for outputs.
5  h_output = cuda.pagelocked_empty(output_nbytes, dtype=np.float32)
6  d_output = cuda.mem_alloc(h_output.nbytes)

```

(continues on next page)

(continued from previous page)

```

7
8 # Transfer data from host to device.
9 cuda.memcpy_htod_async(d_inputs[0], input_a, stream)
10 cuda.memcpy_htod_async(d_inputs[1], input_b, stream)
11 cuda.memcpy_htod_async(d_inputs[2], input_c, stream)
12
13 # Run inference
14 context.execute_async_v2(bindings=[int(d_inp) for d_inp in d_inputs] + [int(d_
    ↳ output)], stream_handle=stream.handle)
15
16 # Synchronize the stream
17 stream.synchronize()

```

After (TensorRT 10.x)

```

1 # Allocate device memory for inputs.
2 d_inputs = [cuda.mem_alloc(input_nbytes) for binding in range(input_num)]
3
4 # Allocate device memory for outputs.
5 h_output = cuda.pagelocked_empty(output_nbytes, dtype=np.float32)
6 d_output = cuda.mem_alloc(h_output.nbytes)
7
8 # Transfer data from host to device.
9 cuda.memcpy_htod_async(d_inputs[0], input_a, stream)
10 cuda.memcpy_htod_async(d_inputs[1], input_b, stream)
11 cuda.memcpy_htod_async(d_inputs[2], input_c, stream)
12
13 # Setup tensor address
14 bindings = [int(d_inputs[i]) for i in range(3)] + [int(d_output)]
15
16 for i in range(engine.num_io_tensors):
17     context.set_tensor_address(engine.get_tensor_name(i), bindings[i])
18
19 # Run inference
20 context.execute_async_v3(stream_handle=stream.handle)
21
22 # Synchronize the stream
23 stream.synchronize()

```

11.2.1. Summary of Changes

- ▶ Added explicit tensor address setup using `set_tensor_address()` with tensor names
- ▶ Changed from `execute_async_v2()` to `execute_async_v3()`
- ▶ The `bindings` parameter is no longer passed to `execute_async_v3()`; tensor addresses must be set beforehand

11.3. Migrating Engine Builds to `build_serialized_network`

The examples below show TensorRT 8.x first, then TensorRT 10.x, for the same engine build path. In TensorRT 10.x, `build_serialized_network()` is the standard build path and is always available. Omit the `build_engine() / serialize()` fallback from the 8.x sample, and check for a `None` return instead of catching `AttributeError`, as shown in the After tab.

Before (TensorRT 8.x)

```

1 engine_bytes = None
2 try:
3     engine_bytes = self.builder.build_serialized_network(self.network, self.
    ↪config)
4 except AttributeError:
5     engine = self.builder.build_engine(self.network, self.config)
6     engine_bytes = engine.serialize()
7     del engine
8 assert engine_bytes

```

After (TensorRT 10.x)

```

1 engine_bytes = self.builder.build_serialized_network(self.network, self.
    ↪config)
2 if engine_bytes is None:
3     log.error("Failed to create engine")
4     sys.exit(1)

```

11.3.1. Summary of Changes

- ▶ The fallback to `build_engine()` and `serialize()` is no longer needed in TensorRT 10.x
- ▶ `build_serialized_network()` is now the standard method and always available
- ▶ Error handling should check for `None` return value instead of catching `AttributeError`

11.4. Python APIs Added in 10.x

The following Python APIs have been added in TensorRT 10.x to support new features and improved functionality.

11.4.1. Types

- ▶ `APILanguage`
- ▶ `ExecutionContextAllocationStrategy`
- ▶ `IGpuAsyncAllocator`

- ▶ InterfaceInfo
- ▶ IPluginResource
- ▶ IPluginV3
- ▶ IStreamReader
- ▶ IVersionedInterface

11.4.2. Methods and Properties

- ▶ ICudaEngine.is_debug_tensor()
- ▶ ICudaEngine.minimum_weight_streaming_budget
- ▶ ICudaEngine.streamable_weights_size
- ▶ ICudaEngine.weight_streaming_budget
- ▶ IExecutionContext.get_debug_listener()
- ▶ IExecutionContext.get_debug_state()
- ▶ IExecutionContext.set_all_tensors_debug_state()
- ▶ IExecutionContext.set_debug_listener()
- ▶ IExecutionContext.set_tensor_debug_state()
- ▶ IExecutionContext.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()

11.5. Removed Python APIs and Replacements

Warning

The APIs listed below have been **removed** in TensorRT 10.x and will cause runtime errors if called. Review each entry for its replacement before upgrading.

Note

The `EngineCapability.kSAFE_DLA` → `EngineCapability.DLA_STANDALONE` rename in the table below is documented for API completeness only. Refer to the **DriveOS 7.2.5 scope** attention notice at the top of this page.

Removed API	Replacement
<code>BuilderFlag.ENABLE_TACTIC_HEURISTIC</code>	Builder optimization level 2.
<code>BuilderFlag.STRICT_TYPES</code>	Use <code>BuilderFlag.DIRECT_IO</code> , <code>BuilderFlag.PREFER_PRECISION_CONSTRAINTS</code> , and <code>BuilderFlag.REJECT_EMPTY_ALGORITHMS</code> .
<code>EngineCapability.DEFAULT</code>	<code>EngineCapability.STANDARD</code>
<code>EngineCapability.kSAFE_DLA</code>	<code>EngineCapability.DLA_STANDALONE</code>
<code>EngineCapability.SAFE_GPU</code>	<code>EngineCapability.SAFETY</code>
<code>IAgorithmIOInfo.tensor_format</code>	Strides, data type, and vectorization information are sufficient to identify tensor formats uniquely.
<code>IBuilder.max_batch_size</code>	Implicit batch support was removed.
<code>IBuilderConfig.max_workspace_size</code>	<code>IBuilderConfig.set_memory_pool_limit()</code> or <code>get_memory_pool_limit()</code> with <code>MemoryPoolType.WORKSPACE</code>
<code>IBuilderConfig.min_timing_iterations</code>	<code>IBuilderConfig.avg_timing_iterations</code>
<code>ICudaEngine.binding_is_input()</code>	<code>ICudaEngine.get_tensor_mode()</code>
<code>ICudaEngine.get_binding_bytes_per_component()</code>	<code>ICudaEngine.get_tensor_bytes_per_component()</code>
<code>ICudaEngine.get_binding_components_per_element()</code>	<code>ICudaEngine.get_tensor_components_per_element()</code>
<code>ICudaEngine.get_binding_dtype()</code>	<code>ICudaEngine.get_tensor_dtype()</code>
<code>ICudaEngine.get_binding_format()</code>	<code>ICudaEngine.get_tensor_format()</code>
<code>ICudaEngine.get_binding_format_desc()</code>	<code>ICudaEngine.get_tensor_format_desc()</code>
<code>ICudaEngine.get_binding_index()</code>	No name-based equivalent replacement.
<code>ICudaEngine.get_binding_name()</code>	No name-based equivalent replacement.
<code>ICudaEngine.get_binding_shape()</code>	<code>ICudaEngine.get_tensor_shape()</code>
<code>ICudaEngine.get_binding_vectorized_dim()</code>	<code>ICudaEngine.get_tensor_vectorized_dim()</code>
<code>ICudaEngine.get_location()</code>	<code>ITensor.location</code>
<code>ICudaEngine.get_profile_shape()</code>	<code>ICudaEngine.get_tensor_profile_shape()</code>
<code>ICudaEngine.get_profile_shape_input()</code>	<code>ICudaEngine.get_tensor_profile_values()</code>
<code>ICudaEngine.has_implicit_batch_dimension()</code>	Implicit batch is no longer supported.
<code>ICudaEngine.is_execution_binding()</code>	No name-based equivalent replacement.
<code>ICudaEngine.is_shape_binding()</code>	<code>ICudaEngine.is_shape_inference_io()</code>

continues on next page

Table 11.1 – continued from previous page

Removed API	Replacement
<code>ICudaEngine.max_batch_size()</code>	Implicit batch is no longer supported.
<code>ICudaEngine.num_bindings()</code>	<code>ICudaEngine.num_io_tensors()</code>
<code>IExecutionContext.get_binding_shape()</code>	<code>IExecutionContext.get_tensor_shape()</code>
<code>IExecutionContext.get_strides()</code>	<code>IExecutionContext.get_tensor_strides()</code>
<code>IExecutionContext.set_binding_shape()</code>	<code>IExecutionContext.set_input_shape()</code>
<code>IFullyConnectedLayer</code>	<code>IMatrixMultiplyLayer</code>
<code>INetworkDefinition.add_convolution()</code>	<code>INetworkDefinition.add_convolution_nd()</code>
<code>INetworkDefinition.add_deconvolution()</code>	<code>INetworkDefinition.add_deconvolution_nd()</code>
<code>INetworkDefinition.add_fully_connected()</code>	<code>INetworkDefinition.add_matrix_multiply()</code>
<code>INetworkDefinition.add_padding()</code>	<code>INetworkDefinition.add_padding_nd()</code>
<code>INetworkDefinition.add_pooling()</code>	<code>INetworkDefinition.add_pooling_nd()</code>
<code>INetworkDefinition.add_rnn_v2()</code>	<code>INetworkDefinition.add_loop()</code>
<code>INetworkDefinition.has_explicit_precision</code>	Explicit precision support was removed in 10.0.
<code>INetworkDefinition.has_implicit_batch_dimension</code>	Implicit batch support was removed.
<code>IRNNv2Layer</code>	<code>ILoop</code>
<code>NetworkDefinitionCreationFlag.EXPLICIT_BATCH</code>	Support was removed in 10.0.
<code>NetworkDefinitionCreationFlag.EXPLICIT_PRECISION</code>	Support was removed in 10.0.
<code>PaddingMode.CAFFE_ROUND_DOWN</code>	Caffe support was removed.
<code>PaddingMode.CAFFE_ROUND_UP</code>	Caffe support was removed.
<code>PreviewFeature.DISABLE_EXTERNAL_TACTIC</code>	External tactics are always disabled for core code.
<code>PreviewFeature.FASTER_DYNAMIC_SHAPES_0</code>	This flag is on by default.
<code>ProfilingVerbosity.DEFAULT</code>	<code>ProfilingVerbosity.LAYER_NAMES_ONLY</code>
<code>ProfilingVerbosity.VERBOSE</code>	<code>ProfilingVerbosity.DETAILED</code>
<code>ResizeMode</code>	Use <code>InterpolationMode</code> . Alias was removed.
<code>SampleMode.DEFAULT</code>	<code>SampleMode.STRICT_BOUNDS</code>
<code>SliceMode</code>	Use <code>SampleMode</code> . Alias was removed.

Chapter 12. Appendix: Migrating C++ Code from TensorRT 8.x to 10.x

This page describes how to update C++ code when you migrate from TensorRT 8.x to 10.x: paired examples for enqueueV3 and the name-based tensor API, guidance on 64-bit dimensions, and lists of C++ APIs added and removed in 10.x.

Attention

DriveOS 7.2.5 scope: DLA-related rename entries documented later on this page (such as `EngineCapability::kSAFE_DLA` → `EngineCapability::kDLA_STANDALONE`) are listed for API completeness only. **DLA is not supported in TensorRT 11.0.1 for NVIDIA DriveOS 7.2.5.** For current DriveOS guidance, refer to [Migrating TensorRT from 10.x to 11.x on NVIDIA DriveOS](#).

12.1. Migrating from enqueueV2 to enqueueV3 (C++)

The examples below show TensorRT 8.x first, then TensorRT 10.x, for the same inference task. In TensorRT 10.x, enqueueV3 replaces enqueueV2: call `setTensorAddress` for each I/O tensor (using names from `getIOTensorName`) before enqueueV3, as shown in the After tab.

Before (TensorRT 8.x)

```
1 // Create RAII buffer manager object.
2 samplesCommon::BufferManager buffers(mEngine);
3
4 auto context = SampleUniquePtr<nvinfer1::IExecutionContext>(mEngine->
5     ↪ createExecutionContext());
6 if (!context)
7 {
8     return false;
9 }
10 // Pick a random digit to try to infer.
```

(continues on next page)

(continued from previous page)

```

11 srand(time(NULL));
12 int32_t const digit = rand() % 10;
13
14 // Read the input data into the managed buffers.
15 // There should be just 1 input tensor.
16 ASSERT(mParams.inputTensorNames.size() == 1);
17
18 if (!processInput(buffers, mParams.inputTensorNames[0], digit))
19 {
20     return false;
21 }
22 // Create a CUDA stream to execute this inference.
23 cudaStream_t stream;
24 CHECK(cudaStreamCreate(&stream));
25
26 // Asynchronously copy data from host input buffers to device input
27 buffers.copyInputToDeviceAsync(stream);
28
29 // Asynchronously enqueue the inference work
30 if (!context->enqueueV2(buffers.getDeviceBindings().data(), stream, nullptr))
31 {
32     return false;
33 }
34 // Asynchronously copy data from device output buffers to host output buffers.
35 buffers.copyOutputToHostAsync(stream);
36
37 // Wait for the work in the stream to complete.
38 CHECK(cudaStreamSynchronize(stream));
39
40 // Release stream.
41 CHECK(cudaStreamDestroy(stream));

```

After (TensorRT 10.x)

```

1 // Create RAII buffer manager object.
2 samplesCommon::BufferManager buffers(mEngine);
3
4 auto context = SampleUniquePtr<nvinfer1::IExecutionContext>(mEngine->
5     ↪createExecutionContext());
6 if (!context)
7 {
8     return false;
9 }
10 for (int32_t i = 0, e = mEngine->getNbIOTensors(); i < e; i++)
11 {
12     auto const name = mEngine->getIOTensorName(i);
13     context->setTensorAddress(name, buffers.getDeviceBuffer(name));
14 }
15
16 // Pick a random digit to try to infer.

```

(continues on next page)

(continued from previous page)

```

17 srand(time(NULL));
18 int32_t const digit = rand() % 10;
19
20 // Read the input data into the managed buffers.
21 // There should be just 1 input tensor.
22 ASSERT(mParams.inputTensorNames.size() == 1);
23
24 if (!processInput(buffers, mParams.inputTensorNames[0], digit))
25 {
26     return false;
27 }
28 // Create a CUDA stream to execute this inference.
29 cudaStream_t stream;
30 CHECK(cudaStreamCreate(&stream));
31
32 // Asynchronously copy data from host input buffers to device input
33 buffers.copyInputToDeviceAsync(stream);
34
35 // Asynchronously enqueue the inference work
36 if (!context->enqueueV3(stream))
37 {
38     return false;
39 }
40
41 // Asynchronously copy data from device output buffers to host output buffers.
42 buffers.copyOutputToHostAsync(stream);
43
44 // Wait for the work in the stream to complete.
45 CHECK(cudaStreamSynchronize(stream));
46
47 // Release stream.
48 CHECK(cudaStreamDestroy(stream));

```

12.1.1. Summary of Changes

- ▶ Added explicit tensor address setup using `setTensorAddress()` with tensor names from `getIOTensorName()`
- ▶ Changed from `enqueueV2()` to `enqueueV3()`
- ▶ The bindings parameter is no longer passed to `enqueueV3()`; tensor addresses must be set beforehand using `setTensorAddress()`

12.2. Migrating Dims and IShapeLayer to 64-Bit Types

Warning

This is a **breaking ABI change**. Code that bitwise copies to or from Dims must be updated for the wider type.

TensorRT 10.x changes the dimension type from `int32_t` to `int64_t`. The dimensions held by Dims changed from `int32_t` to `int64_t`. However, in TensorRT 10.x, 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.

12.3. C++ APIs Added in 10.x

The following C++ APIs have been added in TensorRT 10.x to support new features and improved functionality.

12.3.1. Enums

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

12.3.2. Types

- ▶ `APILanguage`
- ▶ `Dims64`
- ▶ `ExecutionContextAllocationStrategy`
- ▶ `IGpuAsyncAllocator`
- ▶ `InterfaceInfo`
- ▶ `IPluginResource`

- ▶ IPluginV3
- ▶ IStreamReader
- ▶ IVersionedInterface

12.3.3. Methods and Properties

- ▶ `getInferLibBuildVersion`
- ▶ `getInferLibMajorVersion`
- ▶ `getInferLibMinorVersion`
- ▶ `getInferLibPatchVersion`
- ▶ `IBuilderConfig::setMaxNbTactics`
- ▶ `IBuilderConfig::getMaxNbTactics`
- ▶ `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`

12.4. Removed C++ APIs and Replacements

Warning

The APIs listed below have been **removed** in TensorRT 10.x and will cause compilation or linker errors. Review each entry for its replacement before upgrading.

Note

The `EngineCapability::kSAFE_DLA` → `EngineCapability::kDLA_STANDALONE` rename in the table below is documented for API completeness only. Refer to the **DriveOS 7.2.5 scope** attention notice at the top of this page.

Removed API	Replacement
<code>BuilderFlag::kENABLE_TACTIC_HEURISTIC</code>	Builder optimization level 2.
<code>BuilderFlag::kSTRICT_TYPES</code> ¹	Use <code>kREJECT_EMPTY_ALGORITHMS</code> , <code>kDIRECT_IO</code> , and <code>kPREFER_PRECISION_CONSTRAINTS</code> .
<code>EngineCapability::kDEFAULT</code>	<code>EngineCapability::kSTANDARD</code>
<code>EngineCapability::kSAFE_DLA</code>	<code>EngineCapability::kDLA_STANDALONE</code>
<code>EngineCapability::kSAFE_GPU</code>	<code>EngineCapability::kSAFETY</code>
<code>IAlgorithm::getAlgorithmIOInfo()</code>	<code>IAlgorithm::getAlgorithmIOInfoByIndex()</code>
<code>IAlgorithmIOInfo::getTensorFormat()</code>	Strides, data type, and vectorization information are sufficient to identify tensor formats uniquely.
<code>IBuilder::buildEngineWithConfig()</code>	<code>IBuilder::buildSerializedNetwork()</code>
<code>IBuilder::destroy()</code>	<code>delete ObjectName</code>
<code>IBuilder::getMaxBatchSize()</code>	Implicit batch support was removed.
<code>IBuilder::setMaxBatchSize()</code>	Implicit batch support was removed.
<code>IBuilderConfig::destroy()</code>	<code>delete ObjectName</code>
<code>IBuilderConfig::getMaxWorkspaceSize()</code>	<code>IBuilderConfig::getMemoryPoolLimit()</code> with <code>MemoryPoolType::kWORKSPACE</code>
<code>IBuilderConfig::getMinTimingIterations</code>	<code>IBuilderConfig::getAvgTimingIterations()</code>
<code>IBuilderConfig::setMaxWorkspaceSize()</code>	<code>IBuilderConfig::setMemoryPoolLimit()</code> with <code>MemoryPoolType::kWORKSPACE</code>
<code>IBuilderConfig::setMinTimingIterations</code>	<code>IBuilderConfig::setAvgTimingIterations()</code>
<code>IConvolutionLayer::getDilation()</code>	<code>IConvolutionLayer::getDilationNd()</code>
<code>IConvolutionLayer::getKernelSize()</code>	<code>IConvolutionLayer::getKernelSizeNd()</code>

continues on next page

Table 12.1 – continued from previous page

Removed API	Replacement
<code>IC convolutionLayer::getPadding()</code>	<code>IC convolutionLayer::getPaddingNd()</code>
<code>IC convolutionLayer::getStride()</code>	<code>IC convolutionLayer::getStrideNd()</code>
<code>IC convolutionLayer::setDilation()</code>	<code>IC convolutionLayer::setDilationNd()</code>
<code>IC convolutionLayer::setKernelSize()</code>	<code>IC convolutionLayer::setKernelSizeNd()</code>
<code>IC convolutionLayer::setPadding()</code>	<code>IC convolutionLayer::setPaddingNd()</code>
<code>IC convolutionLayer::setStride()</code>	<code>IC convolutionLayer::setStrideNd()</code>
<code>ICudaEngine::bindingIsInput()</code>	<code>ICudaEngine::getTensorIOMode()</code>
<code>ICudaEngine::destroy()</code>	<code>delete ObjectName</code>
<code>ICudaEngine::getBindingBytesPerComponent()</code>	<code>ICudaEngine::getTensorBytesPerComponent()</code>
<code>ICudaEngine::getBindingComponentsPerElement()</code>	<code>ICudaEngine::getTensorComponentsPerElement()</code>
<code>ICudaEngine::getBindingDataType()</code>	<code>ICudaEngine::getTensorDataType()</code>
<code>ICudaEngine::getBindingDimensions()</code>	<code>ICudaEngine::getTensorShape()</code>
<code>ICudaEngine::getBindingFormat()</code>	<code>ICudaEngine::getTensorFormat()</code>
<code>ICudaEngine::getBindingFormatDesc()</code>	<code>ICudaEngine::getTensorFormatDesc()</code>
<code>ICudaEngine::getBindingIndex()</code>	Name-based methods.
<code>ICudaEngine::getBindingName()</code>	Name-based methods.
<code>ICudaEngine::getBindingVectorizedDim()</code>	<code>ICudaEngine::getTensorVectorizedDim()</code>
<code>ICudaEngine::getLocation()</code>	<code>ITensor::getLocation()</code>
<code>ICudaEngine::getMaxBatchSize()</code>	Implicit batch support was removed.
<code>ICudaEngine::getNbBindings()</code>	<code>ICudaEngine::getNbIOTensors()</code>
<code>ICudaEngine::getProfileDimensions()</code>	<code>ICudaEngine::getProfileShape()</code>
<code>ICudaEngine::getProfileShapeValues()</code>	<code>ICudaEngine::getShapeValues()</code>
<code>ICudaEngine::hasImplicitBatchDimension()</code>	Implicit batch support was removed.
<code>ICudaEngine::isExecutionBinding()</code>	No name-based equivalent replacement.
<code>ICudaEngine::isShapeBinding()</code>	<code>ICudaEngine::isShapeInferenceIO()</code>
<code>IDeconvolutionLayer::getKernelSize()</code>	<code>IDeconvolutionLayer::getKernelSizeNd()</code>
<code>IDeconvolutionLayer::getPadding()</code>	<code>IDeconvolutionLayer::getPaddingNd()</code>
<code>IDeconvolutionLayer::getStride()</code>	<code>IDeconvolutionLayer::getStrideNd()</code>
<code>IDeconvolutionLayer::setKernelSize()</code>	<code>IDeconvolutionLayer::setKernelSizeNd()</code>
<code>IDeconvolutionLayer::setPadding()</code>	<code>IDeconvolutionLayer::setPaddingNd()</code>
<code>IDeconvolutionLayer::setStride()</code>	<code>IDeconvolutionLayer::setStrideNd()</code>
<code>IExecutionContext::destroy()</code>	<code>delete ObjectName</code>
<code>IExecutionContext::enqueue()</code>	<code>IExecutionContext::enqueueV3()</code>

continues on next page

Table 12.1 – continued from previous page

Removed API	Replacement
<code>IEExecutionContext::enqueueV2()</code>	<code>IEExecutionContext::enqueueV3()</code>
<code>IEExecutionContext::execute()</code>	<code>IEExecutionContext::executeV2()</code>
<code>IEExecutionContext::getBindingDimension()</code>	<code>IEExecutionContext::getTensorShape()</code>
<code>IEExecutionContext::getShapeBinding()</code>	<code>IEExecutionContext::getTensorAddress()</code> or <code>getOutputTensorAddress()</code>
<code>IEExecutionContext::getStrides()</code>	<code>IEExecutionContext::getTensorStrides()</code>
<code>IEExecutionContext::setBindingDimension()</code>	<code>IEExecutionContext::setInputShape()</code>
<code>IEExecutionContext::setInputShapeBinding()</code>	<code>IEExecutionContext::setInputTensorAddress()</code> or <code>setTensorAddress()</code>
<code>IEExecutionContext::setOptimizationProfile()</code>	<code>IEExecutionContext::setOptimizationProfileAsync()</code>
<code>IFullyConnectedLayer</code>	<code>IMatrixMultiplyLayer</code>
<code>IGpuAllocator::free()</code>	<code>IGpuAllocator::deallocate()</code>
<code>IHostMemory::destroy()</code>	<code>delete ObjectName</code>
<code>INetworkDefinition::addConvolution()</code>	<code>INetworkDefinition::addConvolutionNd()</code>
<code>INetworkDefinition::addDeconvolution()</code>	<code>INetworkDefinition::addDeconvolutionNd()</code>
<code>INetworkDefinition::addFullyConnected()</code>	<code>INetworkDefinition::addMatrixMultiply()</code>
<code>INetworkDefinition::addPadding()</code>	<code>INetworkDefinition::addPaddingNd()</code>
<code>INetworkDefinition::addPooling()</code>	<code>INetworkDefinition::addPoolingNd()</code>
<code>INetworkDefinition::addRNNv2()</code>	<code>INetworkDefinition::addLoop()</code>
<code>INetworkDefinition::destroy()</code>	<code>delete ObjectName</code>
<code>INetworkDefinition::hasExplicitPrecision()</code>	Explicit precision support was removed in 10.0.
<code>INetworkDefinition::hasImplicitBatchDimension()</code>	Implicit batch support was removed.
<code>IOnnxConfig::destroy()</code>	<code>delete ObjectName</code>
<code>IPaddingLayer::getPostPadding()</code>	<code>IPaddingLayer::getPostPaddingNd()</code>
<code>IPaddingLayer::getPrePadding()</code>	<code>IPaddingLayer::getPrePaddingNd()</code>
<code>IPaddingLayer::setPostPadding()</code>	<code>IPaddingLayer::setPostPaddingNd()</code>
<code>IPaddingLayer::setPrePadding()</code>	<code>IPaddingLayer::setPrePaddingNd()</code>
<code>IPoolingLayer::getPadding()</code>	<code>IPoolingLayer::getPaddingNd()</code>
<code>IPoolingLayer::getStride()</code>	<code>IPoolingLayer::getStrideNd()</code>
<code>IPoolingLayer::getWindowSize()</code>	<code>IPoolingLayer::getWindowSizeNd()</code>
<code>IPoolingLayer::setPadding()</code>	<code>IPoolingLayer::setPaddingNd()</code>
<code>IPoolingLayer::setStride()</code>	<code>IPoolingLayer::setStrideNd()</code>
<code>IPoolingLayer::setWindowSize()</code>	<code>IPoolingLayer::setWindowSizeNd()</code>

continues on next page

Table 12.1 – continued from previous page

Removed API	Replacement
<code>IRefitter::destroy()</code>	<code>delete ObjectName</code>
<code>IResizeLayer::getAlignCorners()</code>	<code>IResizeLayer::getAlignCornersNd()</code>
<code>IResizeLayer::setAlignCorners()</code>	<code>IResizeLayer::setAlignCornersNd()</code>
<code>IRuntime::deserializeCudaEngine(void const* blob, std::size_t size, IPluginFactory* pluginFactory)</code>	Use <code>deserializeCudaEngine</code> with two parameters.
<code>IRuntime::destroy()</code>	<code>delete ObjectName</code>
<code>IRNNv2Layer</code>	<code>ILoop</code>
<code>kNV_TENSORRT_VERSION_IMPL²</code>	<code>NV_TENSORRT_VERSION_INT(major, minor, patch)</code> macro: $((\text{major}) * 10000L + (\text{minor}) * 100L + (\text{patch}) * 1L)$
<code>NetworkDefinitionCreationFlag::kEXPLICIT_BATCH</code>	Support was removed in 10.0.
<code>NetworkDefinitionCreationFlag::kEXPLICIT_PRECISION</code>	Support was removed in 10.0.
<code>NV_TENSORRT_SONAME_MAJOR</code>	<code>NV_TENSORRT_MAJOR</code>
<code>NV_TENSORRT_SONAME_MINOR</code>	<code>NV_TENSORRT_MINOR</code>
<code>NV_TENSORRT_SONAME_PATCH</code>	<code>NV_TENSORRT_PATCH</code>
<code>getBuilderSafePluginRegistry()</code>	API will not be implemented.
<code>PaddingMode::kCAFFE_ROUND_DOWN</code>	Caffe support was removed.
<code>PaddingMode::kCAFFE_ROUND_UP</code>	Caffe support was removed.
<code>PreviewFeature::kDISABLE_EXTERNAL_TACTICS</code>	External tactics are always disabled for core code.
<code>PreviewFeature::kFASTER_DYNAMIC_SHAPES</code>	This flag is on by default.
<code>ProfilingVerbosity::kDEFAULT</code>	<code>ProfilingVerbosity::kLAYER_NAMES_ONLY</code>
<code>ProfilingVerbosity::kVERBOSE</code>	<code>ProfilingVerbosity::kDETAILED</code>
<code>ResizeMode</code>	Use <code>InterpolationMode</code> . Alias was removed.
<code>RNNDirection</code>	RNN-related data structures were removed.
<code>RNNGateType</code>	RNN-related data structures were removed.
<code>RNNInputMode</code>	RNN-related data structures were removed.
<code>RNNOperation</code>	RNN-related data structures were removed.
<code>SampleMode::kDEFAULT</code>	<code>SampleMode::kSTRICT_BOUNDS</code>
<code>SliceMode</code>	Use <code>SampleMode</code> . Alias was removed.

¹ When removing enum members (for all enums in this list), we will change the enumeration in the enum to have sequential numbers.

² TensorRT version encoding was changed to accommodate a two-digit minor version.

12.5. Removed C++ Plugins and Replacements

Warning

The plugins listed below have been **removed** in TensorRT 10.x. Using them will cause compilation or linker errors. Review each entry for its replacement before upgrading.

Removed Plugin	Replacement
<code>createAnchorGeneratorPlugin()</code>	<code>GridAnchorPluginCreator::createPlugin()</code>
<code>createBatchedNMSPPlugin()</code>	<code>BatchedNMSPPluginCreator::createPlugin()</code>
<code>createInstanceNormalizationPlugin()</code>	<code>InstanceNormalizationPluginCreator::createPlugin()</code>
<code>createNMSPPlugin()</code>	<code>NMSPPluginCreator::createPlugin()</code>
<code>createNormalizePlugin()</code>	<code>NormalizePluginCreator::createPlugin()</code>
<code>createPriorBoxPlugin()</code>	<code>PriorBoxPluginCreator::createPlugin()</code>
<code>createRegionPlugin()</code>	<code>RegionPluginCreator::createPlugin()</code>
<code>createReorgPlugin()</code>	<code>ReorgPluginCreator::createPlugin()</code>
<code>createRPNROIPlugin()</code>	<code>RPNROIPluginCreator::createPlugin()</code>
<code>createSplitPlugin()</code>	<code>INetworkDefinition::addSlice()</code>
<code>struct Quadruple</code>	Related plugins were removed.

Chapter 13. Appendix: Migrating trtexec Usage from TensorRT 8.x to 10.x

This page describes how to update trtexec usage when migrating from TensorRT 8.x to 10.x: before/after command examples and lists of removed or deprecated options.

13.1. Migrating --workspace and --minTiming Options

The examples below show TensorRT 8.x first, then TensorRT 10.x, for the same trtexec invocation pattern.

Before (TensorRT 8.x)

```
1 trtexec \  
2   --onnx=/path/to/model.onnx \  
3   --saveEngine=/path/to/engine.trt \  
4   --optShapes=input:$INPUT_SHAPE \  
5   \  
6   --workspace=1024 \  
7   --minTiming=1
```

After (TensorRT 10.x)

```
1 trtexec \  
2   --onnx=/path/to/model.onnx \  
3   --saveEngine=/path/to/engine.trt \  
4   --optShapes=input:$INPUT_SHAPE \  
5   \  
6   --memPoolSize=workspace:1024
```

13.1.1. Summary of Changes

- ▶ `--workspace` flag replaced with `--memPoolSize=workspace:<size>`
- ▶ `--minTiming` flag removed (use `--avgTiming` instead)

13.2. Removed trtexec Flags and Replacements

Warning

The flags listed below have been **removed** in TensorRT 10.x. Using them will cause `trtexec` to exit with an error. Review each entry for its replacement before upgrading.

Removed Flag	Replacement
<code>--deploy</code>	TensorRT 10.x does not support Caffe input, UFF input, and implicit batch dimension mode.
<code>--output</code>	TensorRT 10.x does not support Caffe input, UFF input, and implicit batch dimension mode.
<code>--model</code>	TensorRT 10.x does not support Caffe input, UFF input, and implicit batch dimension mode.
<code>--uff</code>	TensorRT 10.x does not support Caffe input, UFF input, and implicit batch dimension mode.
<code>--uffInput</code>	TensorRT 10.x does not support Caffe input, UFF input, and implicit batch dimension mode.
<code>--uffNHWC</code>	TensorRT 10.x does not support Caffe input, UFF input, and implicit batch dimension mode.
<code>--batch</code>	TensorRT 10.x does not support Caffe input, UFF input, and implicit batch dimension mode.
<code>--maxBatch</code>	TensorRT 10.x does not support Caffe input, UFF input, and implicit batch dimension mode.
<code>--minTiming</code>	<code>--avgTiming</code>
<code>--preview=features</code>	<code>disableExternalTacticSourcesForCore0805</code> or <code>fasterDynamicShapes0805</code>
<code>--workspace=N</code>	<code>--memPoolSize=poolspec</code>
<code>--explicitPrecision</code>	Removed (no replacement).
<code>--nativeInstanceNorm</code>	Removed (no replacement).
<code>--heuristic</code>	<code>--builderOptimizationLevel=<N></code> (where <N> can be 0, 1, or 2)
<code>--buildOnly</code>	<code>--skipInference</code>

continues on next page

Table 13.1 – continued from previous page

Removed Flag	Replacement
<code>--nvtxMode</code>	<code>--profilingVerbosity</code>

13.3. Deprecated trtexec Flags and Replacements

The following trtexec flags have been deprecated. Each entry shows the deprecated flag and its replacement.

Deprecated Flag	Replacement
<code>--sparsity=force</code>	Use <code>polygraphy surgeon prune</code> to rewrite the weights to a sparsity pattern and then run <code>--sparsity=enable</code> .
<code>--plugins</code>	<code>--staticPlugins</code>
<code>--preview=profileSharing080</code>	Enabled by default and has no effect.
<code>--profilingVerbosity=default</code>	<code>--profilingVerbosity=layer_names_only</code>
<code>--profilingVerbosity=verbose</code>	<code>--profilingVerbosity=detailed</code>
<code>--streams</code>	<code>--infStreams</code>
<code>--weightless</code>	<code>--stripWeights</code>

Chapter 14. Appendix: Migrating Safety Runtime Code from TensorRT 8.x to 10.x

This page describes how to update safety runtime code when migrating from TensorRT 8.x to 10.x. The safety runtime targets safety-critical applications with clearer error handling and memory management. The numbered walkthrough below pairs 8.x and 10.x C++ snippets for the same tasks, and the final section lists removed safety APIs and replacements.

14.1. Migrating from TensorRT 8.x Safety Runtime to 10.x

The numbered steps below show TensorRT 8.x first, then TensorRT 10.x, for loading the engine, querying I/O tensors, execution context setup, tensor addresses, and inference.

1. Load the engine from the user's local file system into a memory buffer.
2. Create an `IRuntime` object to deserialize the CUDA engine. Use the `IRuntime` object's `deserializeCudaEngine` method to retrieve the `ICudaEngine` object from the memory buffer.

Before (TensorRT 8.x)

```
1 using namespace nvinfer1::safe;
2 auto infer = createInferRuntime(logger);
3 infer->setErrorRecorder(recorder.get());
4 auto engine = infer->deserializeCudaEngine(enginebuffer.data(),
  ↪ enginebufferSize);
```

The `infer` object (an `IRuntime` instance) provides the primary C++ entry point for TensorRT. It deserializes a serialized engine buffer and produces an `ICudaEngine` object.

After (TensorRT 10.x)

```
1 using namespace nvinfer2::safe;
2 ITRTGraph *graph = nullptr;
3 ErrorCode code = createTRTGraph(/* ITRTGraph *& */graph,
4 /*Engine buffer*/ buffer.data(),
```

(continues on next page)

(continued from previous page)

```

5  /*Engine buffer size*/ buffer.size(),
6  /*ISafeRecorder& */ recorder,
7  /*trtManagedScratch*/ true,
8  /*ISafeMemAllocator* */ allocator);

```

ITRTGraph represents a neural network graph. The createTRTGraph function takes the engine data buffer and constructs the graph object from the serialized engine file.

- ▶ recorder – a reference to an ISafeRecorder object (a derived class of IErrorRecorder). A sample implementation is provided.
- ▶ trtManagedScratch – when true (default), the memory allocator (user-supplied or built-in) allocates scratch memory, equivalent to createExecutionContext in TensorRT 8.x. When false (equivalent to createExecutionContextWithoutScratch), the user must supply and manage scratch memory directly.
- ▶ allocator – an optional custom memory allocator implementing the ISafeMemAllocator interface. If nullptr (default), the built-in allocator is used.

3. Get the input/output tensor attributes from the engine.

Before (TensorRT 8.x)

```

1  int32_t nb = engine->getNbIOTensors();
2  char const* inputName = engine->getIOTensorName(inputIndex);
3  char const* outputName = engine->getIOTensorName(outputIndex);
4  Dims inputDims = engine->getTensorShape(inputName);
5  Dims outputDims = engine->getTensorShape(outputName);
6  DataType inputType = engine->getTensorDataType(inputName);
7  DataType outputType = engine->getTensorDataType(outputName);

```

To calculate the tensor volume, you might also need getTensorVectorizedDim, getTensorBytesPerComponent, and getTensorComponentsPerElement.

After (TensorRT 10.x)

```

1  int64_t nb;
2  ErrorCode code = graph->getNbIOTensors(nb);
3  char const* inputName;
4  code = graph->getIOTensorName(inputName, inputIndex);
5  char const* outputName;
6  code = graph->getIOTensorName(outputName, outputIndex);
7  TensorDescriptor inputDescriptor;
8  code = graph->getIOTensorDescriptor(inputDescriptor, inputName);
9  TensorDescriptor outputDescriptor;
10 code = graph->getIOTensorDescriptor(outputDescriptor, outputName);

```

In TensorRT 10.x, the TensorDescriptor convenience class contains all the information needed to construct an IOTensor outside of TensorRT and pass its address back to TensorRT.

```

1  struct TensorDescriptor
2  {
3      //! Name of the IO Tensor.
4      AsciiChar const* tensorName{nullptr};

```

(continues on next page)

(continued from previous page)

```

5      ///! Static shape of the IO Tensor.
6      ///! \note The static shape will depend on the IO Profile selected.
7      PhysicalDims shape{-1, {}};
8      ///! Stride vector for each element of the IO Tensor.
9      PhysicalDims stride{-1, {}};
10     ///! DataType enum for the IO Tensor.
11     ///! \warning The actual type of the tensor data must correspond to this
12     →DataType.
13     DataType dataType{DataType::kFLOAT};
14     ///! The size of the tensor data type in bytes (4 for float and int32, 2
15     →for half, 1 for int8)
16     uint64_t bytesPerComponent{0U};
17     ///! The vector length (in scalars) for a vectorized tensor, 1 if the
18     →tensor is not vectorized.
19     uint64_t componentsPerVector{1U};
20     ///! The dimension index that the tensor is vectorized along, -1 if the
21     →tensor is not vectorized.
22     int64_t vectorizedDim{-1};
23     ///! Total size in bytes for the IO Tensor.
24     uint64_t sizeInBytes{0U};
25     ///! Enum to denote whether the Tensor is for input or output.
26     TensorIOMode ioMode{TensorIOMode::kNONE};
27     ///! Enum to denote whether the Tensor memory is allocated on the GPU,
28     →CPU, or CPU_PINNED.
29     MemoryPlacement memPlacement{MemoryPlacement::kNONE};
30     ///! The order that the dimensions are laid out in memory.
31     PhysicalDims strideOrder{-1, {}};
32     ///! The original user-specified shape of the tensor, provided as a
33     →reference if the tensor is
34     →vectorized. When a tensor is vectorized by size T along some
35     →dimension with size K, the
36     →vectorized dimension is split into two dimensions of sizes ceil(K/
37     →T) and T.
38     Dims userShape{-1, {}};
39 };

```

The `sizeInBytes` field directly provides the memory required for the tensor.

The descriptor captures all properties needed to query, allocate, read, and write IO tensors. The physical tensor shape and layout chosen by the compiler can differ from the original user-specified shape due to padding or tiling for vectorized layouts. For example, a user-specified shape `[1, 10, 3, 4]` with a `TensorFormat` of `NHWC8` tiles by 8 elements along the C dimension, producing the following descriptor fields:

User-specified properties (NHWC8):

- ▶ `userShape`: `[1, 10, 3, 4]` (in NCHW order)
- ▶ `componentsPerVector`: 8
- ▶ `vectorizedDim`: 1 (C dimension)

Compiler-determined layout:

- ▶ `shape`: `[1, 2, 8, 3, 4]` (C dimension split into two dimensions)
- ▶ `stride`: `[192, 8, 1, 64, 16]`

- strideOrder: [4, 1, 0, 3, 2] (tilted C dimensions change fastest)
- 4. Allocate device/host buffers for the input/output tensors of the engine.
- 5. Create an ExecutionContext from the engine with or without scratch memory. If without scratch memory, the user also needs to allocate and set memory for the context.

Before (TensorRT 8.x)

```
1 auto context = engine->createExecutionContext();
```

If without scratchMemory:

```
1 auto context = engine->createExecutionContextWithoutDeviceMemory();
2 size_t size = engine->getDeviceMemorySize();
3 void* mem;
4 cudaError_t code = cudaMalloc(&mem, size);
5 context->setDeviceMemory(mem);
```

After (TensorRT 10.x)

There is no need to call `createExecutionContext` because the `ITRTGraph` object carries its own context.

If `trtManagedScratch` was set to `false` when creating the graph, scratch memory must be supplied manually:

```
1 size_t size;
2 code = graph->getScratchMemorySize(size);
3 void* mem;
4 cudaError_t code = cudaMalloc(&mem, size);
5 code = graph->setScratchMemory(mem);
```

- 6. Preprocess the input data with the host buffer and copy it to the device buffer.
- 7. Set the input/output tensor addresses of the context with the device buffer addresses accordingly.

Before (TensorRT 8.x)

```
1 bool result = context->setInputTensorAddress(inputName, inDataPtr);
2 result = context->setOutputTensorAddress(outputName, outDataPtr);
```

After (TensorRT 10.x)

```
1 TypedArray inputData(inDataPtr, inSize);
2 code = graph->setInputTensorAddress(inputName, inputData);
3 TypedArray outputData(outDataPtr, outSize);
4 graph->setOutputTensorAddress(outputName, outputData);
```

In TensorRT 10.x, IO tensor memory must be strongly typed and wrapped in a `TypedArray` object.

- 8. Call `enqueue` to start inference.

Before (TensorRT 8.x)

```

1 result = context->enqueueV3(stream);
2 cudaStreamSynchronize(stream);

```

After (TensorRT 10.x)

```

1 code = graph->executeAsync(stream);
2 code = graph->sync();

```

In TensorRT 10.x, CUDA stream synchronization is handled by the `sync()` API, which also performs additional finalization.

9. Postprocess the output device buffer to retrieve the inference output.

14.2. Removed Safety C++ APIs and Replacements

Warning

The APIs listed below have been **removed** in TensorRT 10.x and will cause compilation or linker errors. Review each entry for its replacement before upgrading.

The naming pattern follows the general TensorRT migration from binding-based APIs to name-based (tensor) APIs.

Removed API	Replacement
<code>safe::ICudaEngine::bindingIsInput()</code>	<code>safe::ICudaEngine::tensorIOMode()</code>
<code>safe::ICudaEngine::getBindingBytesPerComponent()</code>	<code>safe::ICudaEngine::getTensorBytesPerComponent()</code>
<code>safe::ICudaEngine::getBindingComponentIndex()</code>	<code>safe::ICudaEngine::getTensorComponentsPerElement()</code>
<code>safe::ICudaEngine::getBindingDataType()</code>	<code>safe::ICudaEngine::getTensorDataType()</code>
<code>safe::ICudaEngine::getBindingDimensions()</code>	<code>safe::ICudaEngine::getTensorShape()</code>
<code>safe::ICudaEngine::getBindingIndex()</code>	<code>safe::</code> name-based methods
<code>safe::ICudaEngine::getBindingName()</code>	<code>safe::</code> name-based methods
<code>safe::ICudaEngine::getBindingVectorizedDim()</code>	<code>safe::ICudaEngine::getTensorVectorizedDim()</code>
<code>safe::ICudaEngine::getNbBindings()</code>	<code>safe::ICudaEngine::getNbIOTensors()</code>
<code>safe::ICudaEngine::getTensorFormat()</code>	<code>safe::ICudaEngine::getBindingFormat()</code>
<code>safe::IExecutionContext::enqueueV2()</code>	<code>safe::IExecutionContext::enqueueV3()</code>
<code>safe::IExecutionContext::getStrides()</code>	<code>safe::IExecutionContext::getTensorStrides()</code>

See also

TensorRT C++ API Documentation

Full C++ API documentation.

TensorRT Python API Documentation

Full Python API documentation.

Version Compatibility (in the *TensorRT Developer Guide*)

Engine compatibility across TensorRT versions and hardware.

Troubleshooting (in the *TensorRT Developer Guide*)

Common errors and debugging guidance.

Upgrading TensorRT

Package-level upgrade instructions (pip, Debian, RPM, tar, zip) to complement API-level migration.

TensorRT GitHub: Compare Releases

View API changes between releases using the GitHub compare tool.

Copyright

©2021-2026, NVIDIA Corporation