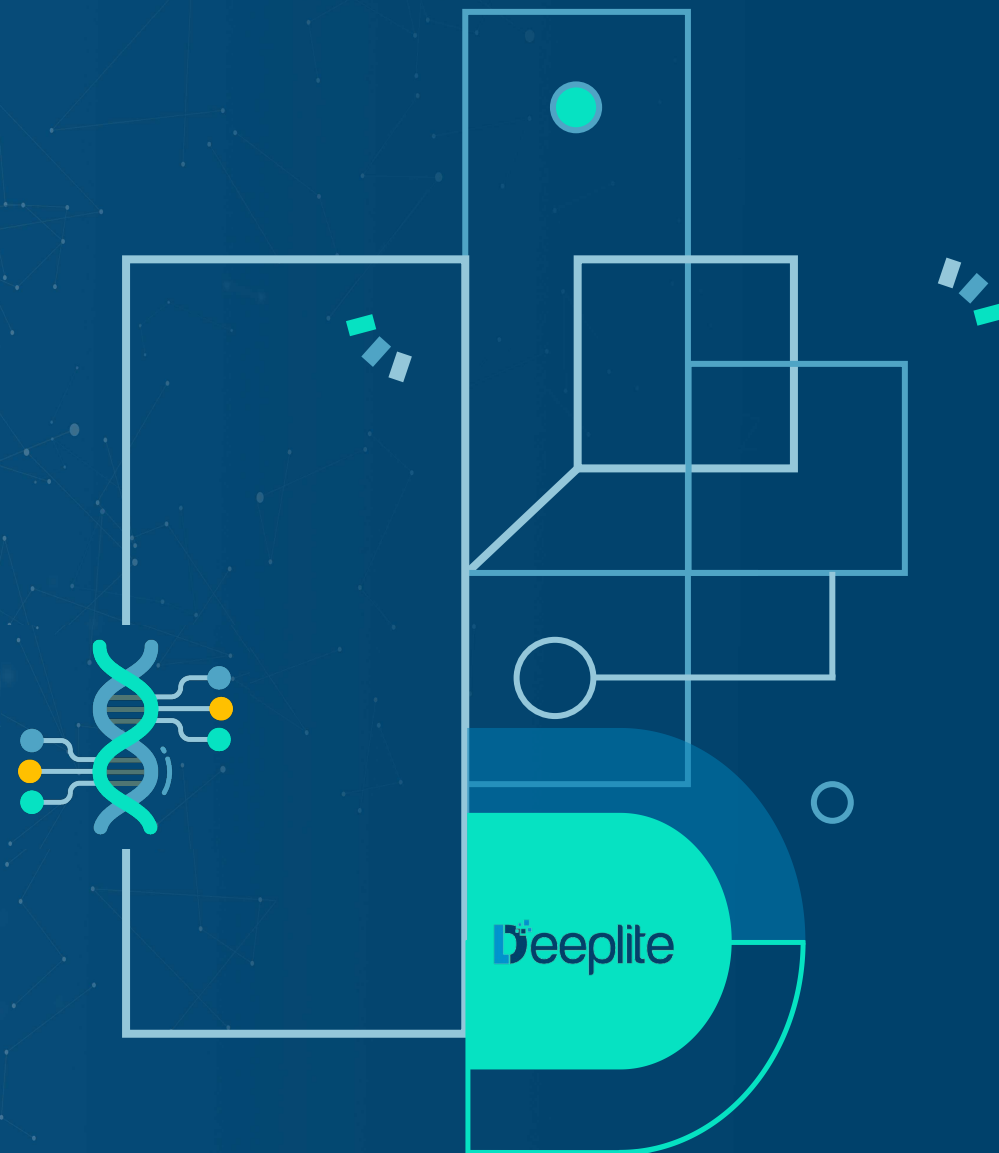


Running 2-bit quantized CNN models on ARM CPUs



Davis Sawyer
Co-founder & CPO



May 4, 2023



Optimize AI for the Edge



Founded 2019



Edge Computing Challenges



High Computational Complexity

Millions of expensive floating-point operations for each input classification are needed.



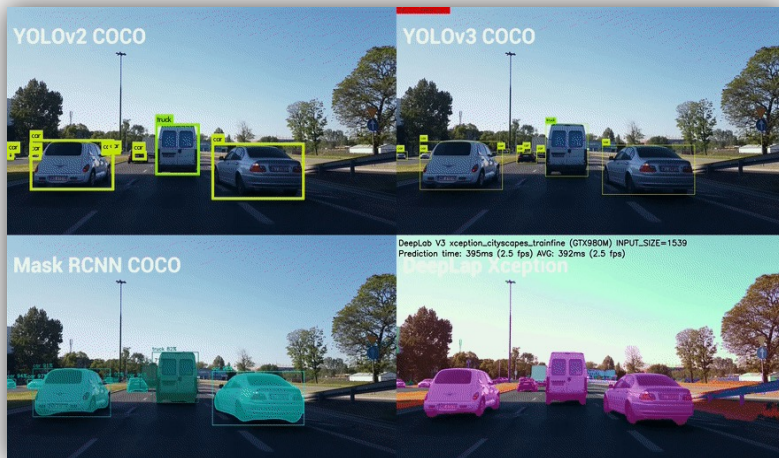
Memory Footprint

Huge amounts of weights and activations with limited on-chip memory and bandwidth.



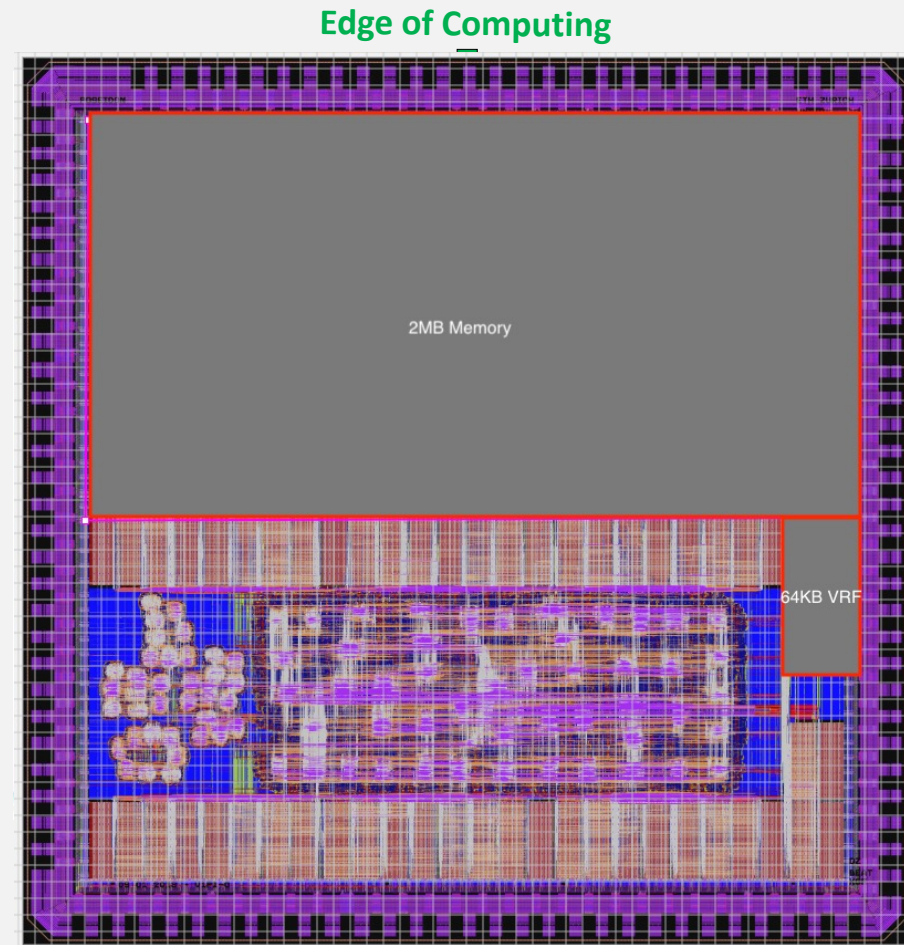
Power consumption

Deep learning requires significant power and can easily consume battery life



In-memory / near memory Computation

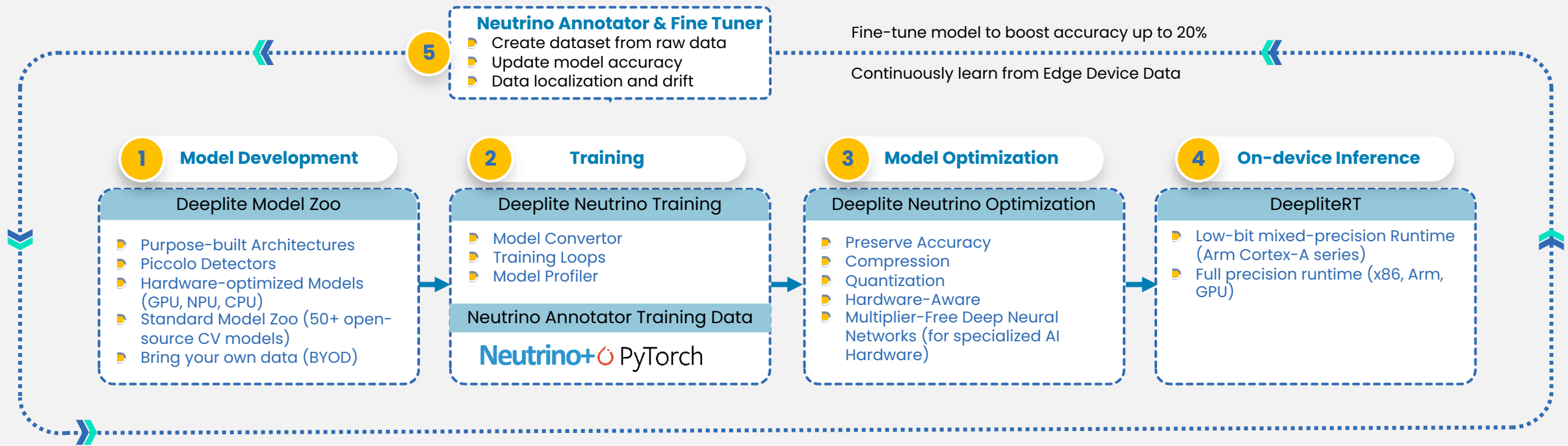
- We must bring computation at the very edge (on-chip, in-memory) Otherwise we lose a lot only on memory bottleneck
- Size is the biggest challenge in this environment



Industrial EdgeAI Challenges

- Big models/datasets 😞
 - Expensive Training Step
- Very sensitive on accuracy drop depends on the use case 😞
- Expecting real time execution on their edge HW 😞
- Generic and automated approaches which work on all use cases 🙏

Deeplite Solution

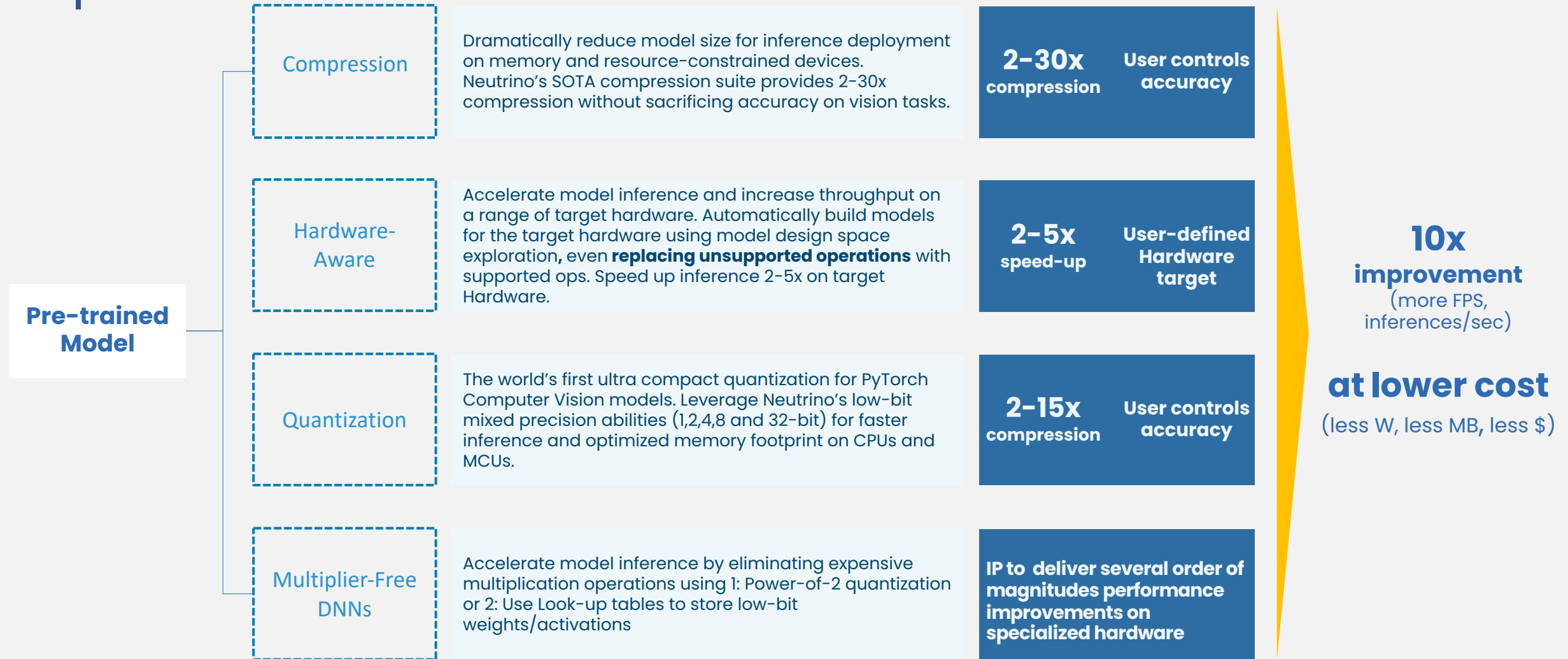


End-to-end Computer Vision Pipeline

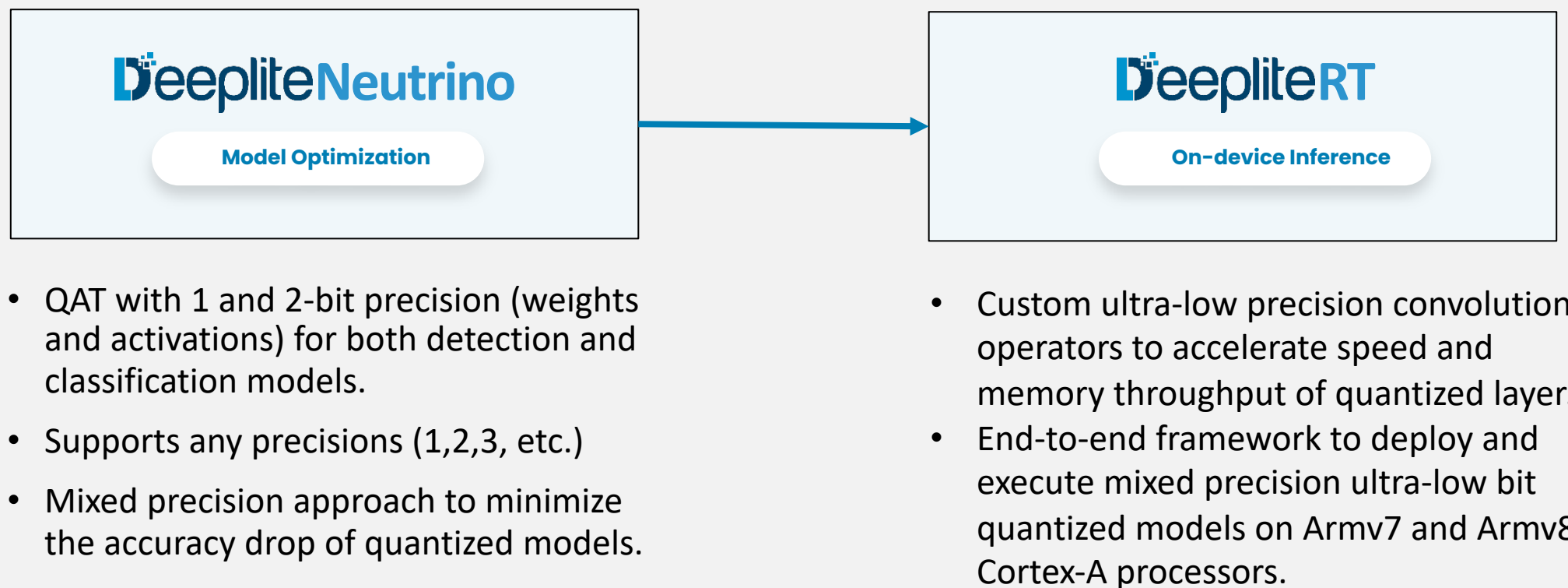
Flexible platform with custom and open-source options

Sophisticated and production-tested model optimization

Neutrino Model Optimization



Our Low-bit Quantization Contribution



Issues with low bit quantization

- High accuracy loss for 2-bit quantization
- Mixed precision quantization
 - Different layers to have different bit precision to maximize accuracy
 - Sensitivity analysis
 - Minimize switching between different precision
 - trade-off between accuracy and speed
 - Automatic way with 2 bit, 8-bit and FP32 layers

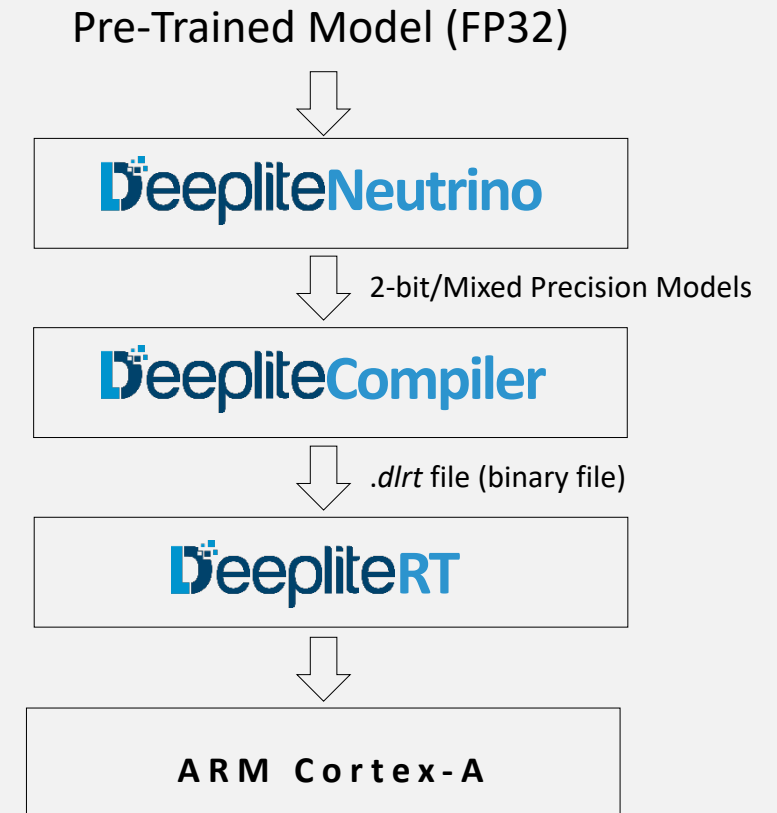
Deeplite Neutrino Quantizer

Sample Applications	Architecture	Deeplite Neutrino Quantization			precision	Accuracy Change	Dataset
		Original Size	Optimized Size	Improvement			
Image classification	ResNet18	42MB	2.9MB	x15	2w/2a	0.00% (Top1)	CIFAR100
	VGG19	76MB	5MB	x15.3	2w/2a	<1.50% (Top1)	CIFAR100
	ResNet18	42MB	2.9MB	x15	2w/1a	<1.00% (Top1)	VWW
	ResNet50	97MB	17MB	x5.6	2w/2a, 8w/8a ¹	~1.50% (Top1)	ImageNet
Object Detection	VGG16_SSD	90MB	5.6MB	x16	2w/2a	<0.01 (mAP)	widerface
	VGG16_SSD	100MB	6.2MB	x16	2w/2a	<0.02 (mAP)	VOC 2012
	Yolo5_6n	7MB	0.95MB	x7	2w/2a, FP32 ¹	<0.02 (mAP)	Custom (Person Detection)

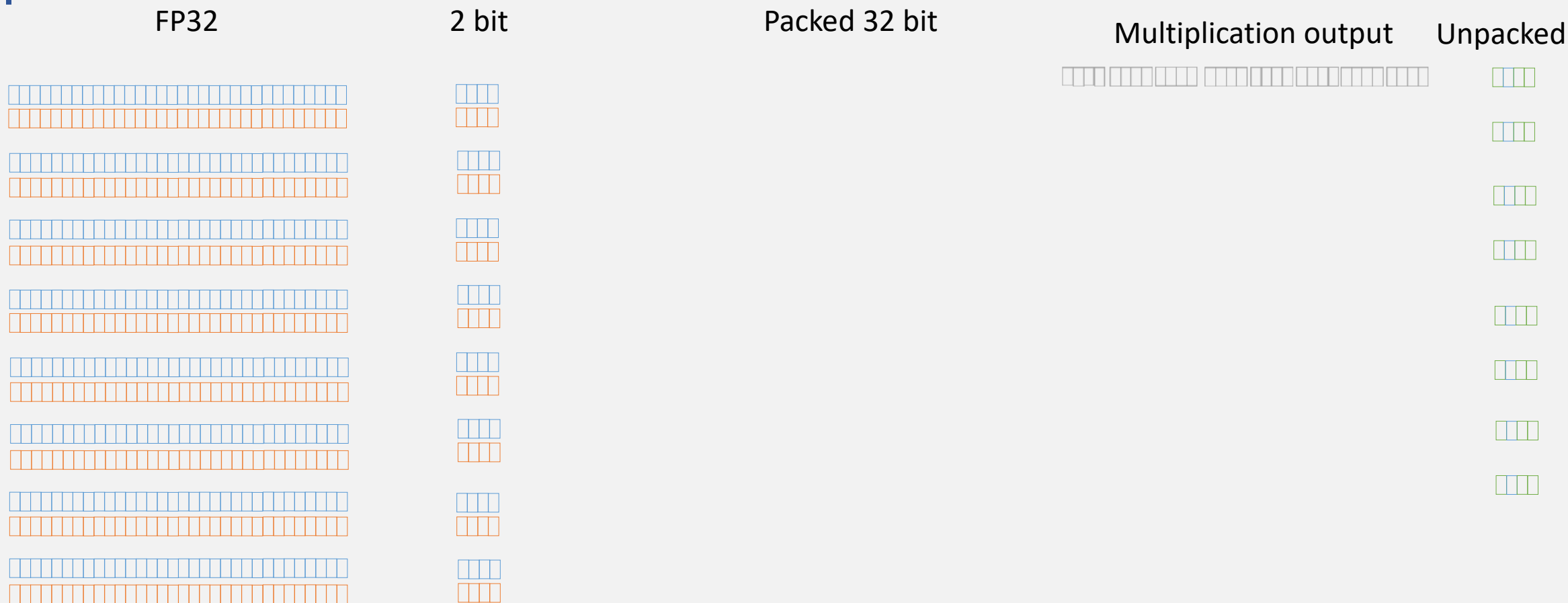
¹Mixed-Precision

DeepliteRT

- Custom ultra-low precision convolution operators to accelerate speed and memory throughput of quantized layers
- End-to-end framework to deploy and execute mixed precision ultra-low bit quantized models on Armv7 and Armv8 Cortex-A processors.



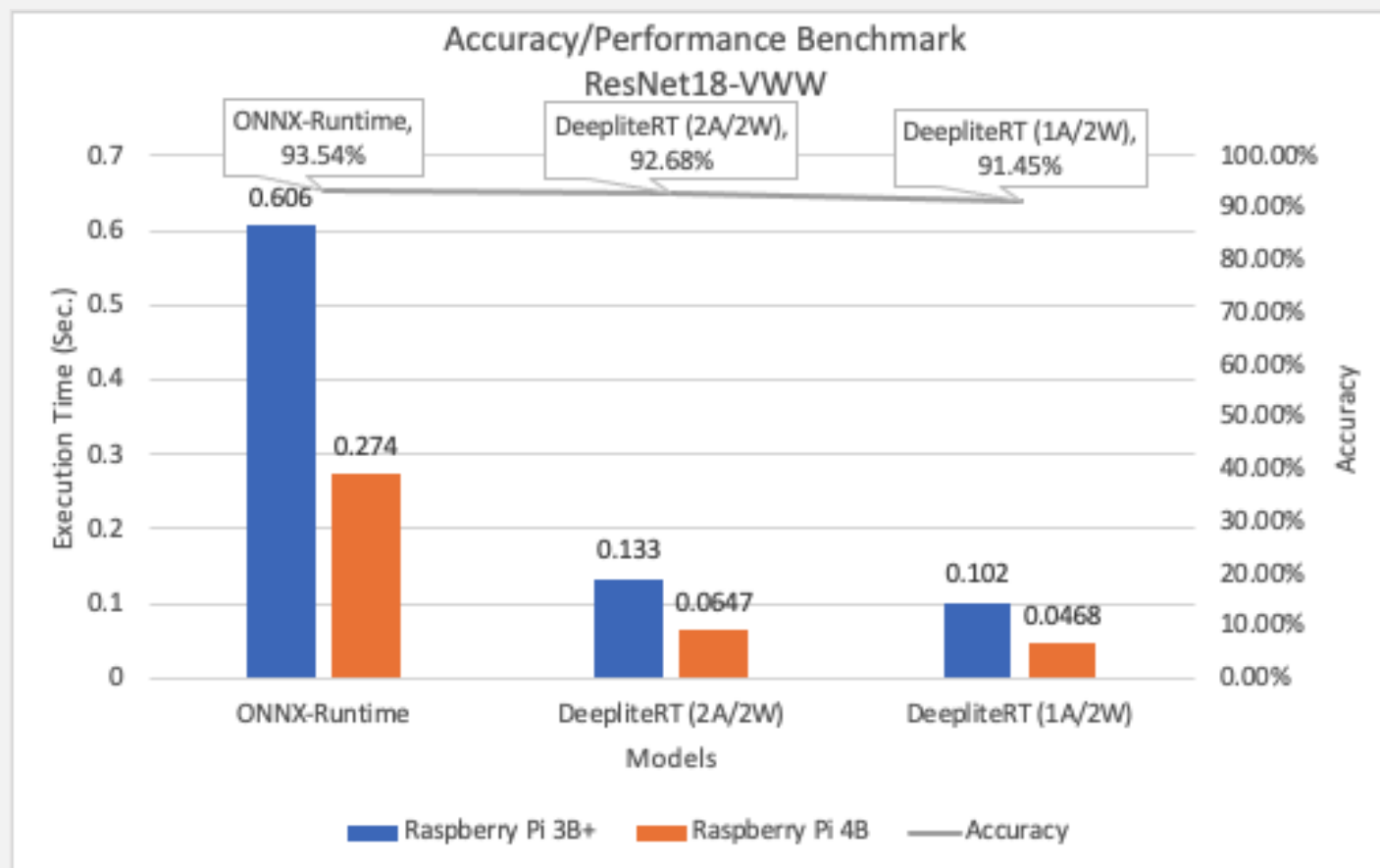
Low Bit Conv2D Implementation



DeepliteRT

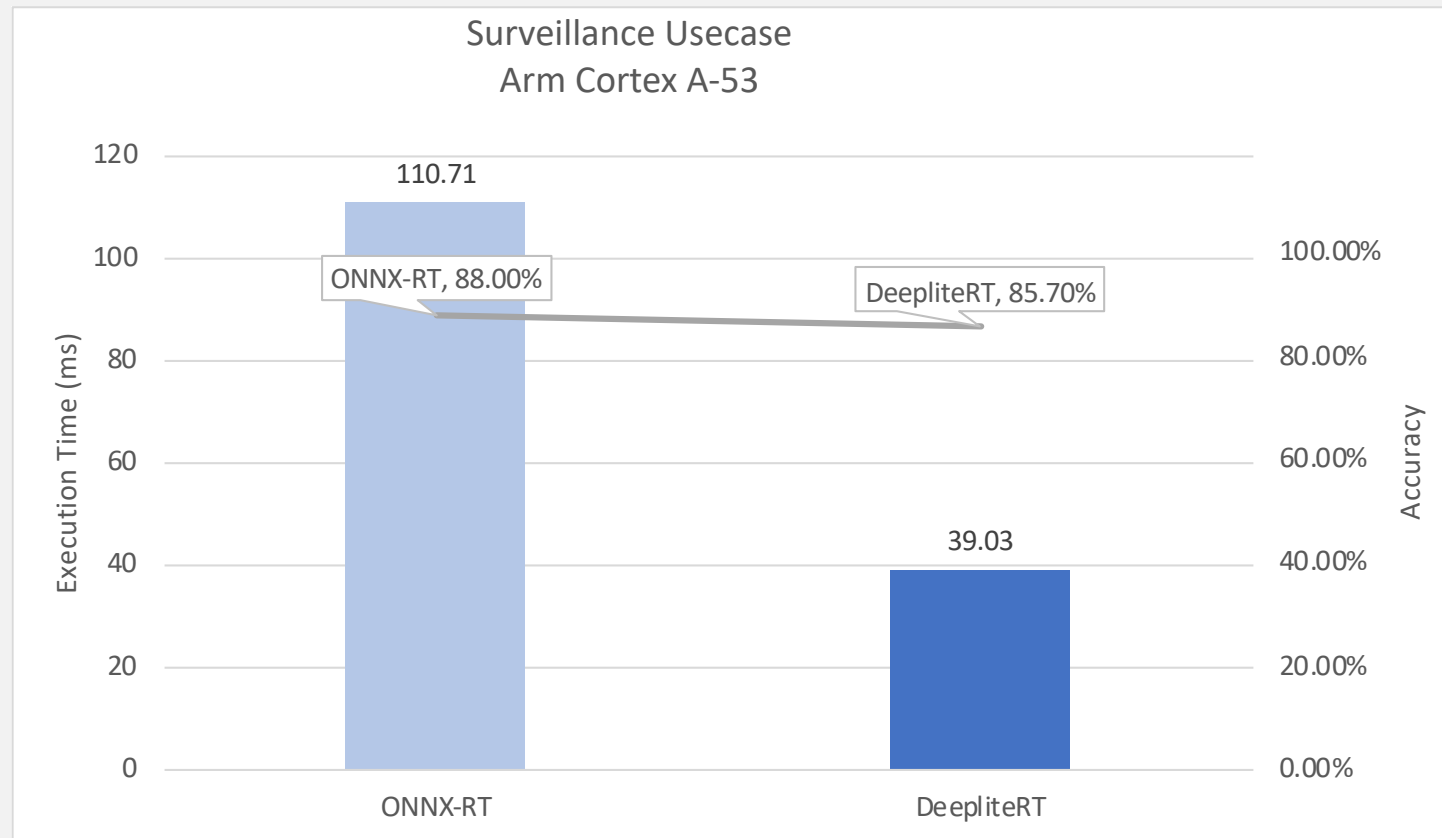
- Uses intrinsic from the Neon vectorized instruction set for both Armv7 and Armv8 architectures to target 32-bit and 64-bit Arm CPU devices.
- Efficient tiling and parallelization schemes used to improve upon the performance of the vectorized kernels.
- On the ResNet18 model running on the low-power Arm Cortex-A53 CPU in the Raspberry Pi 3B+, our overall implementation realizes speedups of up to 2.9x on 2-bit and 4.4x on 1-bit over an optimized floating-point baseline

ResNet18 on VWW



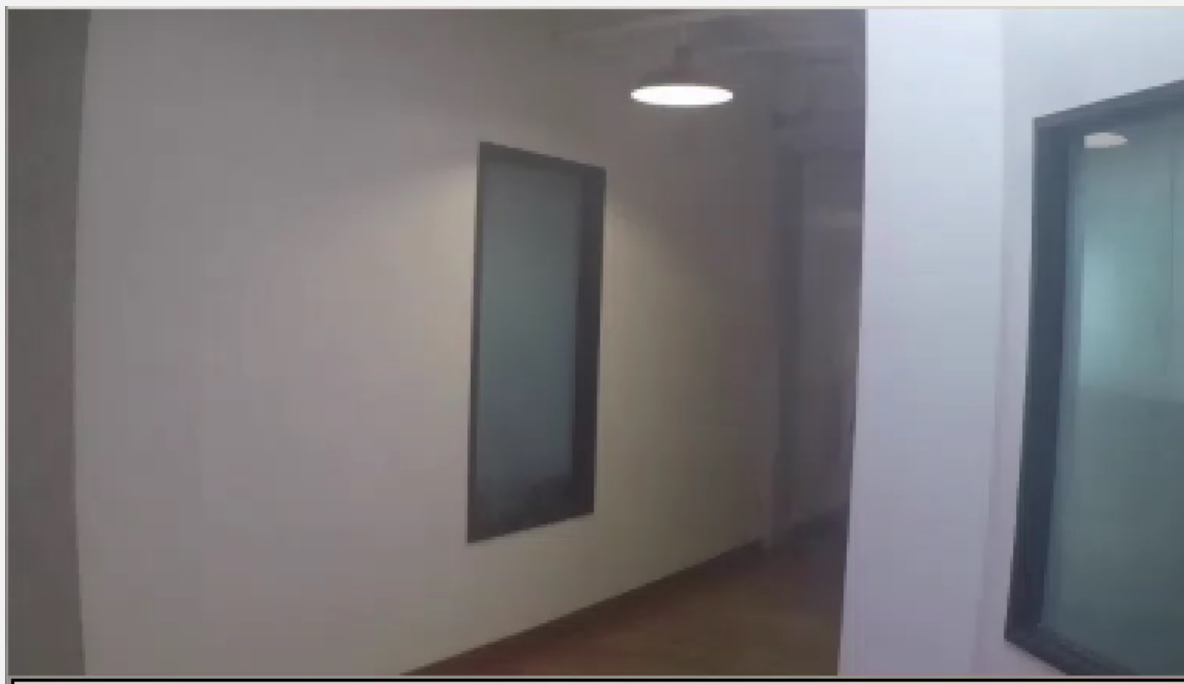
Accuracy/performance benchmark of DeepliteRT on ResNet18 model on VWW dataset (2A/2W- weights and activations quantized to 2 bits, 1A/2W- activations quantized to 1 bit and weights quantized to 2 bits)

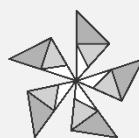
Industrial Use case #1

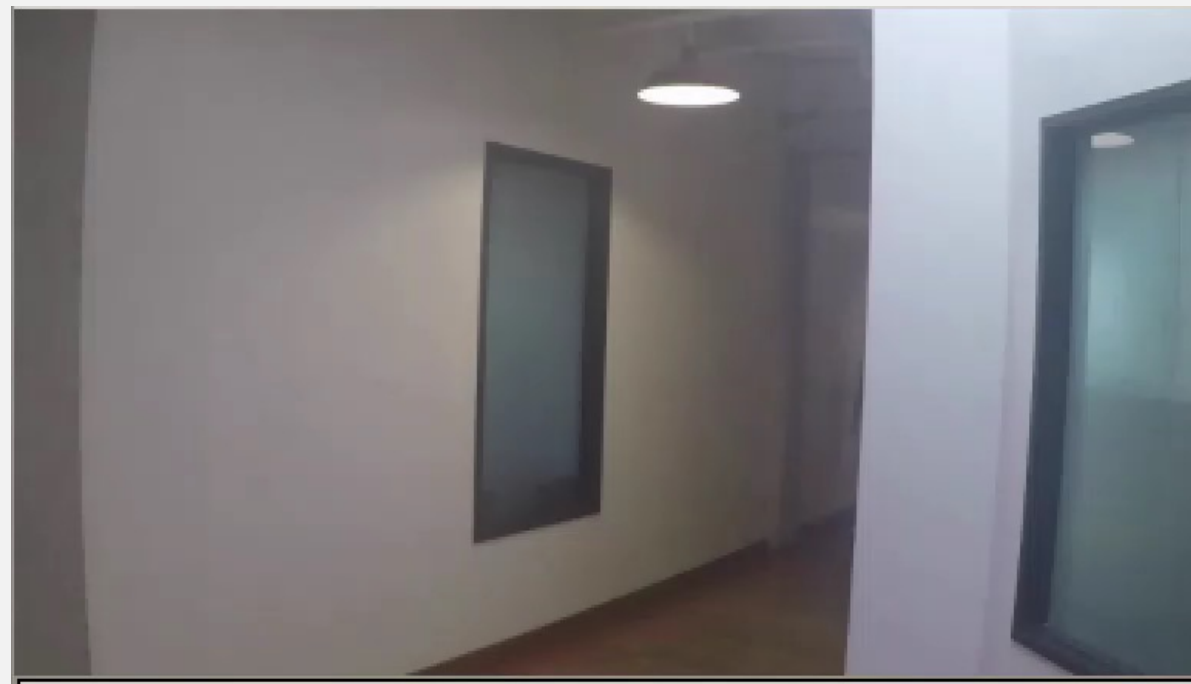


Object detection – Yolo5 based model

DeepliteRT - Detection Performance



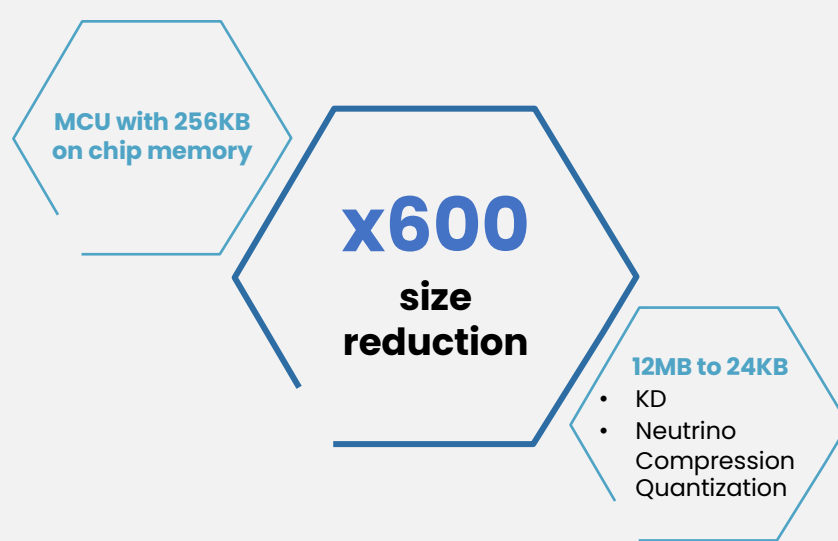
 ONNX
RUNTIME



 DeepliteRT

YOLOv5s COCO – Raspberry Pi 4B (4x Cortex-A72)

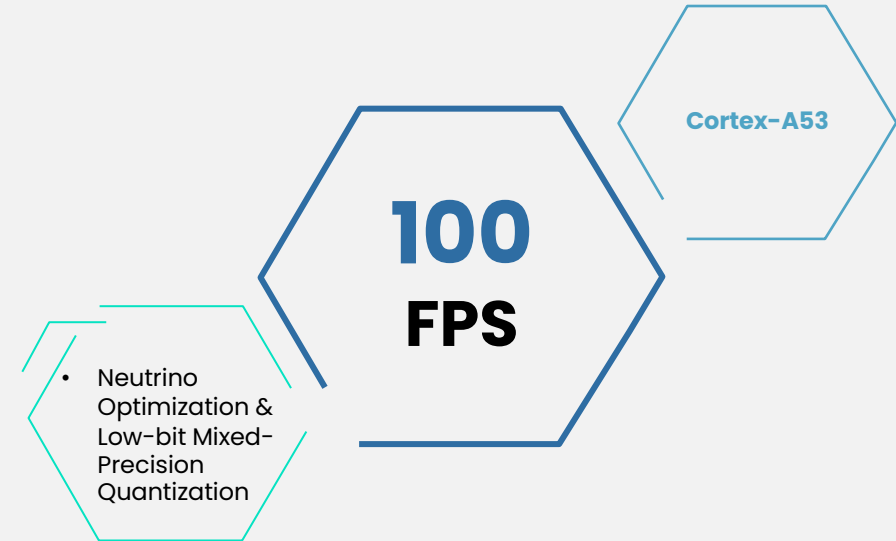
Combining All Techniques



Use Case 1: Person Presence Detection

Baseline MobileNetV1 80% Top1 Accuracy

MLPerf submission is in progress...



Use Case 2: Person Detection

Baseline Yolo5n Model

Conclusion

- 2-bit model running on commodity hardware is presented
- Mixed precision approach is used to minimize the accuracy loss
- Novel method is proposed to run the 2-bit model on Arm Cortex A devices
- Benchmark on classification and Object detection models are presented
- Future work
 - Extend to other hardware
 - Performance Improvement



Thank you

davis@deeplite.ai