

Image Classification						Dev Site
Architecture	Resolution	Dataset	#Params	Quantization	Accuracy	Download
AkidaNet 0.25	160	ImageNet	483K	8	48.50%	Y
				4	41.60%	
AkidaNet 0.5	160	ImageNet	1.4M	8	61.86%	Y
				4	57.93%	
AkidaNet	160	ImageNet	4.4M	8	69.94%	Y
				4	67.23%	
AkidaNet 0.25	224	ImageNet	483K	8	52.39%	Y
				4	46.06%	
AkidaNet 0.5	224	ImageNet	1.4M	8	64.88%	Y
				4	61.47%	
AkidaNet	224	ImageNet	4.4M	8	72.16%	Y
				4	70.11%	
AkidaNet 0.5	224	PlantVillage	1.2M	8	99.61%	Y
				4	98.25%	
AkidaNet 0.25	96	Visual Wake Words	227K	8	87.03%	Y
				4	85.80%	
AkidaNet18	160	ImageNet	2.4M	8	64.77%	Y
AkidaNet18	224	ImageNet	2.4M	8	67.32%	Y
MobileNetV1 0.25	160	ImageNet	469K	8	45.72%	Y
				4	37.51%	
MobileNetV1 0.5	160	ImageNet	1.3M	8	60.27%	Y
				4	54.81%	
MobileNetV1	160	ImageNet	4.2M	8	69.02%	Y
				4	65.28%	
MobileNetV1 0.25	224	ImageNet	469K	8	49.63%	Y
				4	42.08%	
MobileNetV1 0.5	224	ImageNet	1.3M	8	63.65%	Y
				4	59.20%	
MobileNetV1	224	ImageNet	4.2M	8	71.18%	Y
				4	68.52%	
GXNOR	28	MNIST	1.6M	4	98.81%	Y

For 2.0 models, both 8-bit PTQ and 4-bit QAT numbers are given. When not explicitly stated 8-bit PTQ accuracy is given as is. The 4-bit QAT is the same as for 1.0. The quantization scheme indicates both weights and activations bit-width. Weights in the first layer are always quantized to 8-bit. When 'edge' means that the model backbone output is quantized to 1-bit to allow Akida edge learning.

Object Detection						Dev Site
Architecture	Resolution	Dataset	#Params	Quantization	Accuracy	Download
YOLOv2 (AkidaNet 0.5 backbone)	224	PASCAL-VOC 2007	3.6M	8	50.9% / 27.4% / 27.8%	Y
				4	47.2% / 23.3% / 24.7%	
CenterNet (AkidaNet18 backbone)	384	PASCAL-VOC 2007	2.4M	8	70.3% / 47.3% / 43.9%	Y
YOLOv2 (AkidaNet 0.5 backbone)	224	WIDER FACE	3.6M	8	80.19%	Y
				4	78.60%	
Regression						
VGG-like	32	UTKFace (age estimation)	458K	8	6.0304	Y
				4	5.8858	
Face recognition						
AkidaNet 0.5	112×96	CASIA Webface face ID	2.3M	4	69.79%	Y
AkidaNet 0.5	112×96	CASIA Webface face ID	2.3M	8	72.83%	Y
Segmentation						
AkidaUNet 0.5	128	Portrait128	1.1M	8	<u>90.57%</u>	Y
Key Word Spotting						
DS-CNN		Google speech command	23.8K	8	92.87%	Y
	8 bits			4	92.67%	
				4 + edge	90.61%	
Point Cloud Classification						
PointNet++		ModelNet40 3D Point Cloud	605K	8	0.8088	Y
				4	81.56%	

For 2.0 models, both 8-bit PTQ and 4-bit QAT numbers are given. When not explicitly stated 8-bit PTQ accuracy is given as is. The 4-bit QAT is the same as for 1.0. The quantization scheme indicates both weights and activations bit-width. Weights in the first layer are always quantized to 8-bit. When 'edge' means that the model backbone output is quantized to 1-bit to allow Akida edge learning.