

Operator	Input/Data Precision	Weight Precision	Bias/Accumulator Precision	Output Precision	Constraints of params	Source File	Offline analysis File
ArgMax	int8	-	-	int32(value, index)	维数限制: rank>3时shape[0]-1; 输出维度2; workspace 8B.	executor/core/ops/argmax.c; executor/core/ops/venus/	boots/packer/graph_analysis/ops/ArgMax.py
ArgPool2dInt	int8	-	Accumulator int32	int8	NCHW, batch=1, mode: stride 1/2/4, kernel=qp/d; int32累加, 按平台分体使用SM.	executor/core/ops/argpool2dint.c; executor/core/ops/venus/argpool2dint.h	boots/packer/graph_analysis/ops/ArgPool.py
BatchNorm2dInt	int8	Weight int8	Bias/Accumulator int32	int8	NCHW, weight/bias/relu<C; bias scale=qp+qqr; 0<q+qp+qqr<63; 输入输出SM.	executor/core/ops/batchnorm2dint.c; executor/core/ops/	boots/packer/graph_analysis/ops/BatchNorm2dInt.py
BmmInt	int8	int8 int8	Accumulator int32	int8	两输入用dtype, rank 2/3且K固定; 0<q+qp+qp<63; PSRAM输出SM workspace.	executor/core/ops/bmmint.c; executor/core/ops/venus/	boots/packer/graph_analysis/ops/BmmInt.py
Concat	int8 int16 int32	-	-	int8 int16 int32	输入用rank/dtype, fixshape/shape输出; 输出用dtype; 按平台智能选择cpu/psram.	executor/core/ops/concat.c; executor/core/ops/venus/concat.h	boots/packer/graph_analysis/ops/Concat.py
Conv1dInt	int8 int8 packed int8 int8 packed int8 int8 packed	-	Bias/Accumulator int32 Bias/Accumulator int32 Bias/Accumulator int32 Bias/Accumulator int32	int8 int16 int32	N/C/W; dilation=1, stride 1/2/4, kernel= stride; int32 bias/零值; 权重/偏值在SM.	executor/core/ops/conv1dint.c; executor/core/ops/venus/conv1dint.h	boots/packer/graph_analysis/ops/Conv1dInt.py
Conv2dInt	int8 int8 packed	-	Bias/Accumulator int32	int8	NCHW, batch=1; stride 1/2/4, pad=kernel; 0<q+qp+qp<63; int32 bias/零值; 权重/偏值在PSRAM/输出SM.	executor/core/ops/conv2dint.c; executor/core/ops/venus/conv2dint.h	boots/packer/graph_analysis/ops/Conv2dInt.py
ConvTranspose2dInt	int8 int8 packed	-	Bias/Accumulator int32	int8	N/C/W, batch=1, groups=1, dilation=1; stride 1/2/4, output_padding<stride; int32 bias/零值, 权重/偏值在SM.	executor/core/ops/convtranspose2dint.c; executor/core/ops/venus/convtranspose2dint.h	boots/packer/graph_analysis/ops/ConvTranspose2dInt.py
Flatten	Any byte dtype	-	-	-	axis/维度: 仅重新排列shape, 必须保持连续排列; dtype/scale不变.	executor/core/ops/flatten.c; executor/core/ops/venus/	boots/packer/graph_analysis/ops/Flatten.py
Gather	Any byte dtype Any byte dtype	indices int64 indices int32	-	-	indices/0<int32/int64, axis/维度; 输出scale与data/数据; 按元素字节宽度选择.	executor/core/ops/gather.c; executor/core/ops/venus/gather.h	boots/packer/graph_analysis/ops/Gather.py
GLUInt	int8	State int8	Accumulator int32	int8	split轴长度为正偶数; 输入int8; sigmoid及激活shift<43; SM/PSRAM按平台分体.	executor/core/ops/glu.c; executor/core/ops/venus/glu.h	boots/packer/graph_analysis/ops/GLUInt.py
GRUInt	int8	Weights int8	Bias/Accumulator int32	int8	参数: batch=1; 输入权重int8, bias/int32; 输出scale=hidden scale; 状态及workspace 8SM.	executor/core/ops/gru.c; executor/core/ops/venus/gru.h	boots/packer/graph_analysis/ops/GRUInt.py
hqAd	int8 int32 int8 int16 int32	int8 int8 int16 int32	Accumulator int8 Accumulator int32	int8 int16 int32	两输入Zero point=0, 8/q+qp+qp输出; PSRAM输出固定精度SM workspace.	executor/core/ops/hqad.c; executor/core/ops/venus/hqad.h	boots/packer/graph_analysis/ops/hqAd.py
hqDv	int8 int16 int32	Scalar int8 Scalar int16 Scalar int32	Accumulator int8 Accumulator int32	int8 int16 int32	输入Zero point=0; 向量int32; 权重8/q+qp/int8 dtype/数据类型为2/4/8; 0<q+qp+qp<63.	executor/core/ops/hqdv.c; executor/core/ops/venus/hqdv.h	boots/packer/graph_analysis/ops/hqDv.py
hqMl	int8 int16 int32	int8 int16 int32	Accumulator int8 Accumulator int32	int8 int16 int32	输入用dtype, zero point=0; 仍用shape/特征及排列CPU/1; 0<q+qp+qp<63; PSRAM/输出SM.	executor/core/ops/hqml.c; executor/core/ops/venus/hqml.h	boots/packer/graph_analysis/ops/hqMl.py
hqPd	int8	Pads int64, Fill int8	-	int8	int8, rank 1-4; pads/为填充int64, constant/reflect/replicate/按平台智能; PSRAM按精度.	executor/core/ops/hqp.c; executor/core/ops/venus/hqp.h	boots/packer/graph_analysis/ops/hqPd.py
hqSigmoid	int16	-	Accumulator int16	int8	zero point=0; 输出范围int8 Q7; 输入数据类型, shift/按平台智能; PSRAM/按精度SM.	executor/core/ops/hqsigmoid.c; executor/core/ops/venus/hqsigmoid.h	boots/packer/graph_analysis/ops/hqSigmoid.py
hqSub	int8	int8 int8	Accumulator int8	int8	两输入shape int8; qp+qp, qp+qp按平台智能; PSRAM/按精度SM.	executor/core/ops/hqsub.c; executor/core/ops/venus/hqsub.h	boots/packer/graph_analysis/ops/hqSub.py
hqSum	int8	-	Accumulator int32	int8	输出轴长度=0; 输出按精度1用dtype; 0<q+qp+qp<63; 输入输出SM.	executor/core/ops/hqsum.c; executor/core/ops/venus/hqsum.h	boots/packer/graph_analysis/ops/hqSum.py
hqTanh	int16	-	Accumulator int16	int8	输入数据类型Q8; 输出int8 Q7; zero point=0; PSRAM/按平台分体使用SM.	executor/core/ops/hqtanh.c; executor/core/ops/venus/hqtanh.h	boots/packer/graph_analysis/ops/hqTanh.py
hqvar	int8	-	Accumulator int32	int8	int8, 0/1轴前/后维数=2; shift=2*qp+qp按平台智能; 按平台智能使用SM workspace.	executor/core/ops/hqvar.c; executor/core/ops/venus/hqvar.h	boots/packer/graph_analysis/ops/hqvar.py
LayerNormInt	int8	Gamma int8/int16	Bias/Accumulator int32	int8	需要归一化: gamma/bias/大/小/维数, bias int32; 0<q+qp<15 0<q+qp<15+qp+qp<63; 统计缓存SM.	executor/core/ops/layernormint.c; executor/core/ops/venus/layernormint.h	boots/packer/graph_analysis/ops/LayerNormInt.py
LinearInt	int8 int8 packed	int8 packed bias int8/int16/int32 bias int8/int16/int32	Bias/Accumulator int32 Accumulator int32	int8 int16 int32	scale=0-1, 输入rank 1, 平台/精度; weight/rank2, bias/按精度/精度; shift/按精度/精度; pack/DMA/按精度/精度.	executor/core/ops/linear.c; executor/core/ops/venus/linear.h	boots/packer/graph_analysis/ops/LinearInt.py
LqSoftmaxInt	int8	-	Accumulator int32	int8	按精度/精度, 长度1, 2048; zero point=0; 输入rank Q15, 输出按Q15/精度; workspace/按精度/精度.	executor/core/ops/lqsoftmaxint.c; executor/core/ops/venus/lqsoftmaxint.h	boots/packer/graph_analysis/ops/LqSoftmaxInt.py
LSTMInt	int8	Weights int8	Bias/Accumulator int32; Cell int32 Q15	int8	参数/参数, rank3, batch=1; 输入权重int8, bias int32, cell int32 Q15; 状态/状态 workspace 8SM.	executor/core/ops/lstm.c; executor/core/ops/venus/lstm.h	boots/packer/graph_analysis/ops/LstmInt.py
MaxPool	int8	-	-	int8	NCHW int8, 半/全; stride 1/2/4, kernel= stride; 无tensor/PSRAM按平台分体使用SM.	executor/core/ops/maxpool.c; executor/core/ops/venus/maxpool.h	boots/packer/graph_analysis/ops/Pool.py
PreLU	int8 int16 int32	Slope int32/attr Slope int32/attr Slope int32/attr	-	int8 int16 int32	输入输出shape/dtype/精度/精度; slope/post_shift/0..63 0/1<63<63.	executor/core/ops/prelu.c; executor/core/ops/venus/prelu.h	boots/packer/graph_analysis/ops/activation/PreLU.py
ReLU	int8 int16 int32 int16 int32 int16 int32 int16 int32 int16 int32	-	-	int8 int16 int32 int16 int32 int16 int32 int16 int32 int16 int32	0<q+qp+qp<63; SM/按精度/精度; 任一输入PSRAM/按精度/精度; workspace/按精度/精度.	executor/core/ops/relu.c; executor/core/ops/venus/relu.h	boots/packer/graph_analysis/ops/ReLU.py
Requant	int8 int8 int8	-	-	int8 int16 int32	输入int8, zero point=0; scale/2/2/2, 精度/精度/精度; 按精度/精度/精度; workspace/按精度/精度.	executor/core/ops/requant.c; executor/core/ops/venus/requant.h	boots/packer/graph_analysis/ops/Requant.py
Reshape	Any byte dtype	-	-	-	需要shape/可重排轴/按精度/精度; dtype/scale/不变; 按精度/精度.	executor/core/ops/reshape.c; executor/core/ops/venus/reshape.h	boots/packer/graph_analysis/ops/Reshape.py
ShuffleChannel	-	-	-	-	平台/精度: 按精度/精度/精度/精度/精度/精度; 按精度/精度/精度/精度/精度/精度; 按精度/精度/精度/精度/精度/精度.	executor/core/ops/shufflechannel.c; executor/core/ops/venus/shufflechannel.h	boots/packer/graph_analysis/ops/unsupported/ShuffleChannel.py 未支持, 由 operator_int.h 未支持
Slice	Any byte dtype	starts/ends/axis/steps int32	-	-	按精度/精度/精度/精度/精度/精度; axis/精度, steps=1, dtype/scale/不变; 按精度/精度/精度/精度/精度/精度.	executor/core/ops/slice.c; executor/core/ops/venus/slice.h	boots/packer/graph_analysis/ops/Slice.py
SoftmaxInt	int8	-	Accumulator int32	int8	按精度/精度, 长度1, 2048; zero point=0; 输入rank Q15, 输出按Q15/精度; workspace/按精度/精度.	executor/core/ops/softmaxint.c; executor/core/ops/venus/softmaxint.h	boots/packer/graph_analysis/ops/SoftmaxInt.py
Split	int8 int16 int32	-	-	int8 int16 int32	axis/精度, split/精度/精度/精度/精度/精度; 按精度/精度/精度/精度/精度/精度.	executor/core/ops/split.c; executor/core/ops/venus/split.h	boots/packer/graph_analysis/ops/Split.py
Tile	int8	Repeats int64	-	int8	repeats/精度/精度/精度/精度/精度/精度; 按精度/精度/精度/精度/精度/精度.	executor/core/ops/tile.c; executor/core/ops/venus/tile.h	boots/packer/graph_analysis/ops/Tile.py
topk	int8	index offset int64	-	int8	按精度/精度, max_num=1, shape/0-1; offset/精度/精度; 按精度/精度/精度/精度/精度/精度.	executor/core/ops/topk.c; executor/core/ops/venus/topk.h	boots/packer/graph_analysis/ops/topk.py
topk2	int16(value, index)	-	-	int16(value, index)	输入rank 3且/精度; 按精度/精度, max_num=1; 输入按精度/精度/精度/精度/精度/精度.	executor/core/ops/topk2.c; executor/core/ops/venus/topk2.h	boots/packer/graph_analysis/ops/topk2.py
Transpose	int8 int16 int32	-	-	int8 int16 int32	axis/精度/精度/精度/精度/精度/精度; 按精度/精度/精度/精度/精度/精度.	executor/core/ops/transpose.c; executor/core/ops/venus/transpose.h	boots/packer/graph_analysis/ops/Transpose.py
Quant	Float32	-	-	int8	按精度/精度, 输入按精度/精度, zero point=0, scale/精度/精度/精度/精度/精度.	executor/core/ops/quant.c; executor/core/ops/venus/quant.h	boots/packer/graph_analysis/ops/Quant.py
Dequant	int8(int8/int32)	-	-	Float32	按精度/精度, 输入按精度/精度, zero point=0, scale/精度/精度/精度/精度/精度.	executor/core/ops/dequant.c; executor/core/ops/venus/dequant.h	boots/packer/graph_analysis/ops/Dequant.py

ShuffleChannel					不支持：没有源时和未导入分析源文件，仅缺少执行注册及分析源导入。源通模型不可用。	executor/core/ops/shufflechannel.c; executor/core/ops/venus/shufflechannel.h	unsupported: ShuffleChannel.py 未导入，且 operator_sict.h 未注册
Slice	Any byte dtype	start%ends/axes/step int32/int64		Same as input	步数为零名需要整型int32/int64; axis合法, step=1; dtype/scale不变。早选精片PSRAH使用SM。	executor/core/ops/slice.c; executor/core/ops/arc/slice.h	boots/packer/graph_analysis/ops/Slice.py
SoftmaxInt	int8	-	Accumulator int32	int8	仅最后轴，长度1..2048; zero point=0; 输入为Q25，输出为Q15量化; workspace为int32指数/步和缓冲。	executor/core/ops/softmaxint.c; executor/core/ops/arc/softmaxint.h	boots/packer/graph_analysis/ops/Softmaxint.py
StarifyFFint	int8	Weights int4/int8 * Mask	BlockAccumulator int32	int8	runtime-only: 无普通函数analysis/设备入口; 精度布局, 精度及workspace由源构造ARCS/设备保证。	executor/core/ops/starifyffint.c; executor/core/ops/arc/starifyffint.h	runtime-only: 无普通函数 analysis
Split	int8	-	-	int8	axis合法，各split长度之和等于轴长; 输出dtype/scale同输入; PSRAH设备使用SM/DMA。	executor/core/ops/split.c; executor/core/ops/arc/split.h	boots/packer/graph_analysis/ops/Split.py
	int16	-	-	int16			
	int32	-	-	int32			
Tile	int8	Repeats int64		int8	repeats为变量int64且长度=rank，各值>0; dtype/scale不变; 输出重复块数量使用SM。	executor/core/ops/tile.c; executor/core/ops/arc/tile.h	boots/packer/graph_analysis/ops/Tile.py
topN	int8	Index offset int64		int32(value,index)	仅最后轴，max_num=1, shape[0]=1; offset为int64; 输出按索引对，workspace按平台限定，workspace 16B。	executor/core/ops/topn.c; executor/core/ops/arc/topn.h	boots/packer/graph_analysis/ops/topN.py
topN2	int32(value,index)	-		int32(value,index)	输入rank-3维数据，仅最后轴，max_num=1; 输入输出scale相同并位于SM; workspace按平台限定，workspace 16B。	executor/core/ops/topn2.c; executor/core/ops/arc/topn2.h	boots/packer/graph_analysis/ops/topN2.py
Transpose	int8	-		Same as input	rank/perf满足平台自命名; dtype/scale不变; 矩阵需要为PSRAH设备由SM workspace的表。	executor/core/ops/transpose.c; executor/core/ops/arc/transpose.h	boots/packer/graph_analysis/ops/Transpose.py
	int32	-		int32			
Quant	Float32	-		int8	输入边界转换: 输入输出shape相同, zero point=0, scale为2之幂; 当前有限输出仅int8，不能原样复用。	executor/core/ops/quant.c; executor/core/comm/sdfile.c	boots/packer/graph_analysis/ops/Quant.py
Dequant	int8[1/16/32]/int32	-		Float32	输入边界转换: 输入输出shape相同, zero point=0, scale指数为整数且在[0..30]; Float32输出不得继续进入普通计算算子。	executor/core/ops/dequant.c; executor/core/comm/sdfile.c	boots/packer/graph_analysis/ops/Dequant.py

Operator	Input Data Precision	Weight Precision	Bias/Accumulator Precision	Output Precision	constraints of params	Source File	Offline analysis File
ArgMax	int8	-	-	int32 (cpu) or int8 (psa)	从最后轴: rank>1时shape[0]=1; 输出直接2; workspace 8B.	executor/core/ops/argmax.c; executor/core/ops/venusA/argmax.h	tools/packer/graph_analysis/ops/ArgMax.py
	int16	-	-	int32 (cpu) or int8 (psa)			
	int32	-	-	int32 (cpu) or int8 (psa)			
AvgPool2dInt	int8	-	Accumulator int32	int8	NCHW, batch=1, ceil; stride 1/2/4, kernel= stride; int32累加, 跨平台分块使用SM.	executor/core/ops/avgpool2dint.c; executor/core/ops/venusA/avgpool2dint.h	tools/packer/graph_analysis/ops/Pool.py
BatchNorm2dInt	int8	Weight int8	Bias/Accumulator int32	int8	NCHW, weight/bias长度=C, bias scale=qx+qx; 0<=qx+qx-qx<=63; 输入输出SM.	executor/core/ops/batchnorm2dint.c; executor/core/ops/venusA/batchnorm2dint.h	tools/packer/graph_analysis/ops/batchnorm2dint.py
BmmInt	int8	RHS same dtype	Accumulator int32	int8			
	int8	RHS same dtype	Accumulator int32	int16			
	int8	RHS same dtype	Accumulator int32	int32			
	int16	RHS same dtype	Accumulator int32	int8			
	int16	RHS same dtype	Accumulator int32	int16	跨输入用dtype, rank 2/3且K取1; 0<=qx+qx-qx<=63; PSAAH输出至SM workspace.	executor/core/ops/bmmint.c; executor/core/ops/venusA/bmmint.h	tools/packer/graph_analysis/ops/bmmint.py
	int16	RHS same dtype	Accumulator int32	int32			
	int32	RHS same dtype	Accumulator int32	int8			
	int32	RHS same dtype	Accumulator int32	int16			
	int32	RHS same dtype	Accumulator int32	int32			
Concat	int8	-	-	int8			
	int16	-	-	int16	输入用rank/dtype, 跨axis shape拼接; 输出用dtype; 跨平台需对scale移位, PSAAH需SM.	executor/core/ops/concat.c; executor/core/ops/venusA/concat.h	tools/packer/graph_analysis/ops/Concat.py
	int32	-	-	int32			
Conv1dInt	int8	int4 packed	Bias/Accumulator int32	int8			
	int8	int4 packed	Bias/Accumulator int32	int16			
	int8	int4 packed	Bias/Accumulator int32	int32			
	int8	int8 packed	Bias/Accumulator int32	int8	N,C,W; dilation=1, stride 1/2/4, kernel= stride; int32 bias/零加; 权重直接连接并跨SM按量分块.	executor/core/ops/conv1dint.c; executor/core/ops/venusA/conv1dint.h	tools/packer/graph_analysis/ops/venusA/Conv1dint.py
	int8	int8 packed	Bias/Accumulator int32	int16			
	int8	int8 packed	Bias/Accumulator int32	int32			
Conv2dInt	int8	int4 packed	Bias/Accumulator int32	int8	NCHW, batch=1; stride 1/2/4, pad=kernel; 0<=qx+qx-qx<=63; int32 bias/零加, 权重直接及PSAAH分块输入SM.	executor/core/ops/conv2dint.c; executor/core/ops/venusA/conv2dint.h	tools/packer/graph_analysis/ops/venusA/Conv2dint.py
	int8	int8 packed	Bias/Accumulator int32	int8			
ConvTranspose2dInt	int8	int4 packed	Bias/Accumulator int32	int8	NCHW, batch=1, group=1, dilation=1; stride 1/2/4, output_padding< stride; int32 零加, 权重直接并跨SM.	executor/core/ops/convtranspose2dint.c; executor/core/ops/venusA/convtranspose2dint.h	tools/packer/graph_analysis/ops/venusA/ConvTranspose2dint.py
	int8	int8 packed	Bias/Accumulator int32	int8			
Flatten	Any byte dtype	-	-	Same as input	axis合法; 仅重解释shape, 必需对选择拷贝; dtype/scale不变.	executor/core/ops/flatten.c; executor/core/ops/venusA/flatten.h	tools/packer/graph_analysis/ops/Flatten.py
Gather	Any byte dtype	Indices int64	-	Same as data			
	Any byte dtype	Indices int32	-	Same as data	indices仅int32/int64, axis合法; 输出scale与data相同; 跨元素字节宽度适配.	executor/core/ops/gather.c; executor/core/ops/venusA/gather.h	tools/packer/graph_analysis/ops/Gather.py
Gelu	int8	-	Accumulator int32 Q27	int8			
	int16	-	Accumulator int32 Q27	int8	仅VenusA; 输入SM, 输出int8; scale数据为-3.90146(-36.57); int8的输入元素, int16/32 跨轴均需workspace.	executor/core/ops/gelu.c; executor/core/ops/venusA/gelu.h	tools/packer/graph_analysis/ops/Gelu.py
	int32	-	Accumulator int32 Q27	int8			
GLuInt	int8	Gate int8	Accumulator int32	int8			
	int8	Gate int8	Accumulator int32	int16	split轴长度为正偶数; 输入int8; sigmoid及乘法shift<=63; SM/PSAAH按平台分块.	executor/core/ops/gluInt.c; executor/core/ops/venusA/gluInt.h	tools/packer/graph_analysis/ops/GLUInt.py
	int8	Gate int8	Accumulator int32	int32			
GRuInt	int8	Weights int8	Bias/Accumulator int32	int8	批量, batch=1; 输入权重int8, bias/门零加int32; 输出scale=hidden scale; 状态及workspace至SM.	executor/core/ops/gruInt.c; executor/core/ops/venusA/gruInt.h	tools/packer/graph_analysis/ops/GRUInt.py
iqAdd	int8	RHS same dtype/same shape	Accumulator same dtype	int8			
	int16	RHS same dtype/same shape	Accumulator same dtype	int16	跨shape且zero point=0; 仅qx-qy-qz-qz移位; PSAAH或重指定使用SM workspace, VenusA或按权重直接, 仅用shape/int8/16/32/64/128输入输出.	executor/core/ops/iqadd.c; executor/core/ops/venusA/iqadd.h	tools/packer/graph_analysis/ops/iqAdd.py
	int32	RHS same dtype/same shape	Accumulator same dtype	int32			
iqDiv	int32	RHS int32 vector	Accumulator int32	int32	输入输出SM; 向量仅int32; 权重ps/rms/归约dtype直接为int32不需; 0<=qx-qy-qz-qz<=63.	executor/core/ops/iqdiv.c; executor/core/ops/venusA/iqdiv.h	tools/packer/graph_analysis/ops/iqDiv.py
	int32	Scalar int32	-	int32			
iqMul	int8	RHS int8 same/scalar	Accumulator int32	int8			
	int16	RHS int16 same/scalar	Accumulator int32	int8			
	int16	RHS int16 same/scalar	Accumulator int32	int16	输入用dtype, zero point=0; 仅用shape/权重直接并跨轴C117"需; 0<=qx+qx-qx-qx<=63; PSAAH分块至SM, VenusA或按权重直接NCHW/NC11; 不支持任意rank"需.	executor/core/ops/iqmul.c; executor/core/ops/venusA/iqmul.h	tools/packer/graph_analysis/ops/iqMul.py
	int32	RHS int32 same/scalar	Accumulator int32	int32			
iqPad	int8	Pads int64, Fill int8(0)	-	int8	int8, rank 1-4; pads为成对int64, constant/reflect/replicate按平台限制; PSAAH按量分块.	executor/core/ops/iqpad.c; executor/core/ops/venusA/iqpad.h	tools/packer/graph_analysis/ops/iqPad.py
iqSigmoid	int8	-	Accumulator int32	int8			
	int16	-	Accumulator int32	int8	zero point=0, 输出固定int8 Q7; 输入权重Q26, shift按平台限制; PSAAH分块使用SM.	executor/core/ops/iqsigmoid.c; executor/core/ops/venusA/iqsigmoid.h	tools/packer/graph_analysis/ops/iqSigmoid.py
	int32	-	Accumulator int32	int8			
iqSub	int8	RHS int8	Accumulator int8	int8	跨输入用shape int8; qx-qy, qy-qz在平台范围; PSAAH或重指定至SM.	executor/core/ops/iqsub.c; executor/core/ops/venusA/iqsub.h	tools/packer/graph_analysis/ops/iqSub.py
iqSum	int8	-	Accumulator int32	int8			
	int16	-	Accumulator int32	int16	仅最后轴直接长度=0; 输出按直接1并开dtype; 0<=qx-qx-qx<=63; 输入输出SM.	executor/core/ops/iqsum.c; executor/core/ops/venusA/iqsum.h	tools/packer/graph_analysis/ops/iqSum.py
	int32	-	Accumulator int32	int32			
iqTanh	int32 Q27	-	Accumulator int32	int8 Q7	输入权重Q26, 输出int8 Q7; zero point=0; PSAAH按平台分块至SM.	executor/core/ops/iqtanh.c; executor/core/ops/venusA/iqtanh.h	tools/packer/graph_analysis/ops/iqTanh.py
iqVar	int8	-	Accumulator int32	int8	int8, 仅用最后轴数据第二轴; shift=2*qx-qy-qz按平台限制; 按直接量使用SM workspace.	executor/core/ops/iqvar.c; executor/core/ops/venusA/iqvar.h	tools/packer/graph_analysis/ops/iqVar.py

LayerNormInt	int8	Gamma int8	BackAccumulator int32	int8	模型1.0推理一体化: gammaBack大小调整, bias int32; 0<=qx<=15且0<=15+qx-qp<=63; 统计缓存在SM,	executor/core/ops/layernormint.c; executor/core/ops/venusA/layernormint.h	tools/packer/graph_analysis/ops/LayerNormInt.py
LinearInt	int8	int4 packed	BackAccumulator int32	int8	transB=1, 输入rank 1, 平台1层: weightTrans2, bias匹配输出宽度; shift按输出宽度预移; axis0DMA前置输入SM, 显式组合: (0,8w4)>0; (0,8w6)>0; (0,16w12); (0,32w32)>32; (16,16w16)>8; (16,32w32)>8; (32,32w32)>8; (32,8w8)>8,	executor/core/ops/linearint.c; executor/core/ops/venusA/linearint.h	tools/packer/graph_analysis/ops/LinearInt.py
	int8/int16	Weight same dtype	BackAccumulator int32	int8/int16/int32			
	int8	int32	BackAccumulator int32	int32			
	int32	int32	BackAccumulator int32	int8/int32			
LinearInt	int32	int8	BackAccumulator int32	int8			
	int32	int8	BackAccumulator int32	int8			
	int32	int8	BackAccumulator int32	int8			
	int32	int8	BackAccumulator int32	int8			
LogSoftmaxInt	int8	-	Accumulator int32	int8	仅最后轴, 长度1, 2048; zero point=0; 输入int32并前置scale量校准; workspace按轴共享,	executor/core/ops/logsoftmaxint.c; executor/core/ops/venusA/logsoftmaxint.h	tools/packer/graph_analysis/ops/LogSoftmaxInt.py
LSTMInt	int8	Weights int8	BackAccumulator int32; Cell int32 Q15	int8	模型串并, rank3, batch=1; 输入权重int8, bias int32, cell int32 Q15; 状态与workspace在SM,	executor/core/ops/lstmint.c; executor/core/ops/venusA/lstmint.h	tools/packer/graph_analysis/ops/LstmInt.py
MaxPool	int8	-	-	int8	NCWH int8, 零cell; stride 1/2/4, kernel>=stride; 大tensor/PSRAM扁平自分块到SM,	executor/core/ops/maxpool.c; executor/core/ops/venusA/maxpool.h	tools/packer/graph_analysis/ops/Pool.py
PReLU	int8	Slope int32(attr)	-	int8	输入输出shape/dtype都符合SM; slopepoint_shift0_63匹配<=63,	executor/core/ops/prelu.c; executor/core/ops/venusA/prelu.h	tools/packer/graph_analysis/ops/activations.py
	int16	Slope int32(attr)	-	int16			
	int32	Slope int32(attr)	-	int32			
ReLU	int8	-	-	int8	0<=qp-qx<=63; SM断接按表支持; 任一输入PSRAM时仅int8->int8, workspace最多64KB,	executor/core/ops/relu.c; executor/core/ops/venusA/relu.h	
	int8	-	-	int16			
	int8	-	-	int32			
	int16	-	-	int8			
	int16	-	-	int16			
	int16	-	-	int32			
	int32	-	-	int8			
	int32	-	-	int16			
	int32	-	-	int32			
ReLUx	int8	-	-	int8	输入输出SM; 输出int8; threshold输入int8, shift 0_63,	executor/core/ops/relu.x.c; executor/core/ops/venusA/relu.x.h	
Requant	int8	-	-	int8	输入int8, zero point=0; scale为2次幂, 位宽/shift在平台限制; 用位宽输入输出带SM,	executor/core/ops/requant.c; executor/core/ops/venusA/requant.h	tools/packer/graph_analysis/ops/Requant.py
	int8	-	-	int16			
	int8	-	-	int32			
Reshape	Any byte dtype	-	-	Same as input	调整shape可连续映射元素数不变; dtype/scale不变, 通常到6,	executor/core/ops/reshape.c; executor/core/ops/venusA/reshape.h	tools/packer/graph_analysis/ops/Reshape.py
Slice	Any byte dtype	start/end/axes/shape int32/int64	-	Same as input	参数为元素数量int32/int64; axis合法, step=1; dtype/scale不变, 平选用PSRAM使用SM,	executor/core/ops/slice.c; executor/core/ops/venusA/slice.h	tools/packer/graph_analysis/ops/Slice.py
SoftmaxInt	int8	-	Accumulator int32	int8	仅最后轴, 长度1, 2048; zero point=0; 输入int32, 输出由Q15量校准; workspace按int32数据深度或减半,	executor/core/ops/softmaxint.c; executor/core/ops/venusA/softmaxint.h	tools/packer/graph_analysis/ops/SoftmaxInt.py
	int8	-	Accumulator int32	int16			
	int8	-	Accumulator int32	int32			
	int16	-	Accumulator int32	int8			
	int16	-	Accumulator int32	int16			
	int16	-	Accumulator int32	int32			
	int32	-	Accumulator int32	int8			
	int32	-	Accumulator int32	int16			
Split	int8	-	-	int8	axis合法, 自split长度之和等于轴长; 输出dtype/scale同输入; PSRAM断接按表支持DMA,	executor/core/ops/split.c; executor/core/ops/venusA/split.h	tools/packer/graph_analysis/ops/Split.py
	int16	-	-	int16			
	int32	-	-	int32			
Swish	int8	-	Accumulator int32 Q27	int8	仅VenusA; 输入输出SM, 输出int8; scale值数x(3,90)split_36,57; int8的轴0元素, int16/32的轴1元素workspace,	executor/core/ops/swish.c; executor/core/ops/venusA/swish.h	tools/packer/graph_analysis/ops/Swish.py
	int16	-	Accumulator int32 Q27	int8			
	int32	-	Accumulator int32 Q27	int8			
Title	int8	Repeats int64	-	int8	repeats为变量int64且长度=rank, 各值>0; dtype/scale不变; 输出重复按原量使用SM,	executor/core/ops/title.c; executor/core/ops/venusA/title.h	tools/packer/graph_analysis/ops/Title.py
topN	int8	index offset int64	-	int32(value,index)	仅最后轴, max_num=1, shape[0]=1; offset为int64; 输出带索引, workspace按平台固定, workspace 16B,	executor/core/ops/topn.c; executor/core/ops/venusA/topn.h	tools/packer/graph_analysis/ops/topN.py
topK2	int32(value,index)	-	-	int32(value,index)	输入rank3且带轴2, 仅最后轴, max_num=1; 输入输出scale相同并位于SM; workspace按平台固定, workspace 16B,	executor/core/ops/topn2.c; executor/core/ops/venusA/topn2.h	tools/packer/graph_analysis/ops/topK2.py
Transpose	int8	-	-	Same as input	rank/permut满足平台自命名; dtype/scale不变; 转得数量及PSRAM断接由SM workspace的数,	executor/core/ops/transpose.c; executor/core/ops/venusA/transpose.h	tools/packer/graph_analysis/ops/Transpose.py
	int16	-	-	int16			
	int32	-	-	int32			
Quant	Float32	-	-	int8	模型边界转码; 输入输出shape相同, zero point=0, scale为2次方幂; 当前有转码输出int8, 不能断接复用,	executor/core/ops/quant.c; executor/core/comm/pttic.c	tools/packer/graph_analysis/ops/Quant.py
Dequant	int8/int8/int32	-	-	Float32	模型边界转码; 输入输出shape相同, zero point=0, scale值为整数且在0.30; Float32输出不得继续进入后续计算算子,	executor/core/comm/pttic.c	tools/packer/graph_analysis/ops/Dequant.py