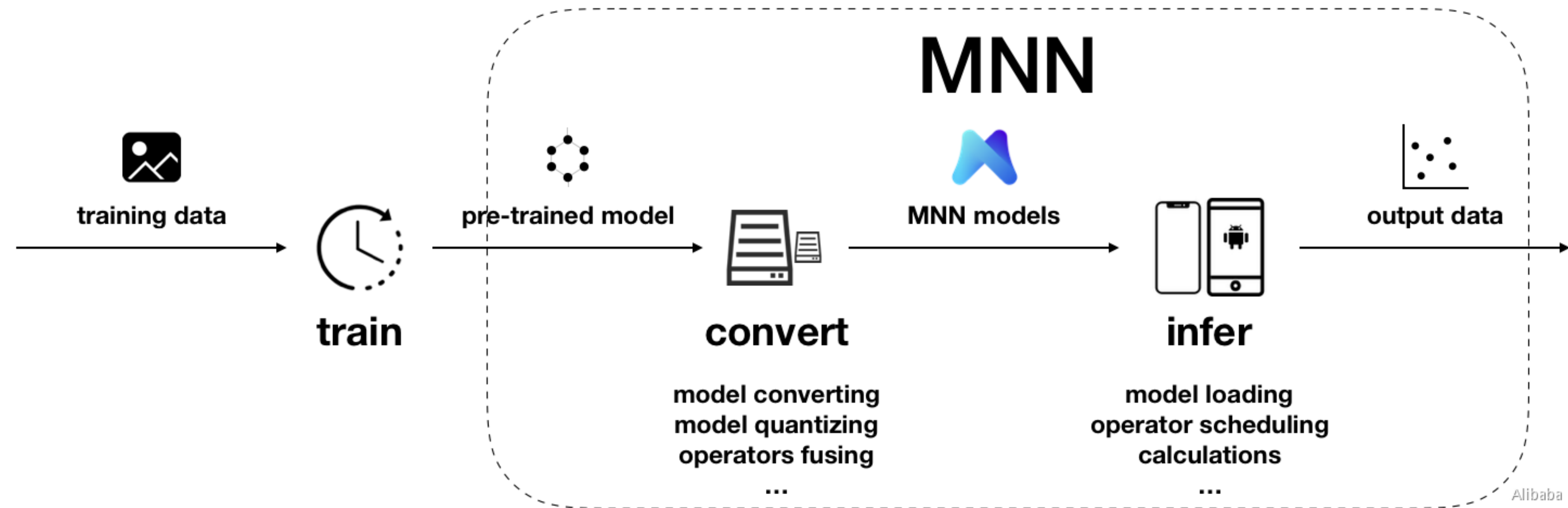


MNN's Solution of Heterogeneous Computing - Geometry Compute

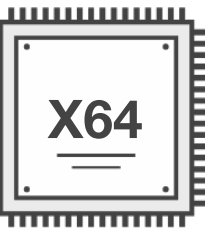
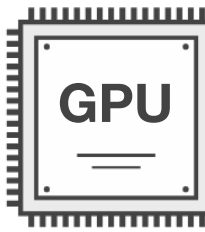
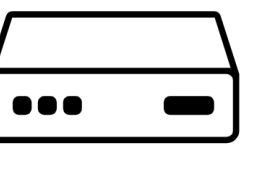
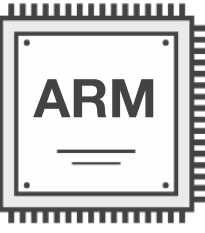
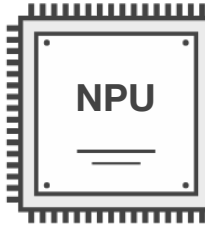
Jiang Xiaotang



MNN - A Universal and Fast Inference Engine

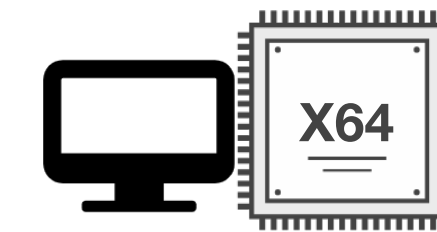


MNN Support:

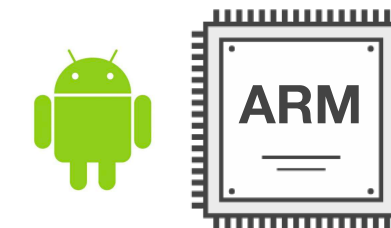
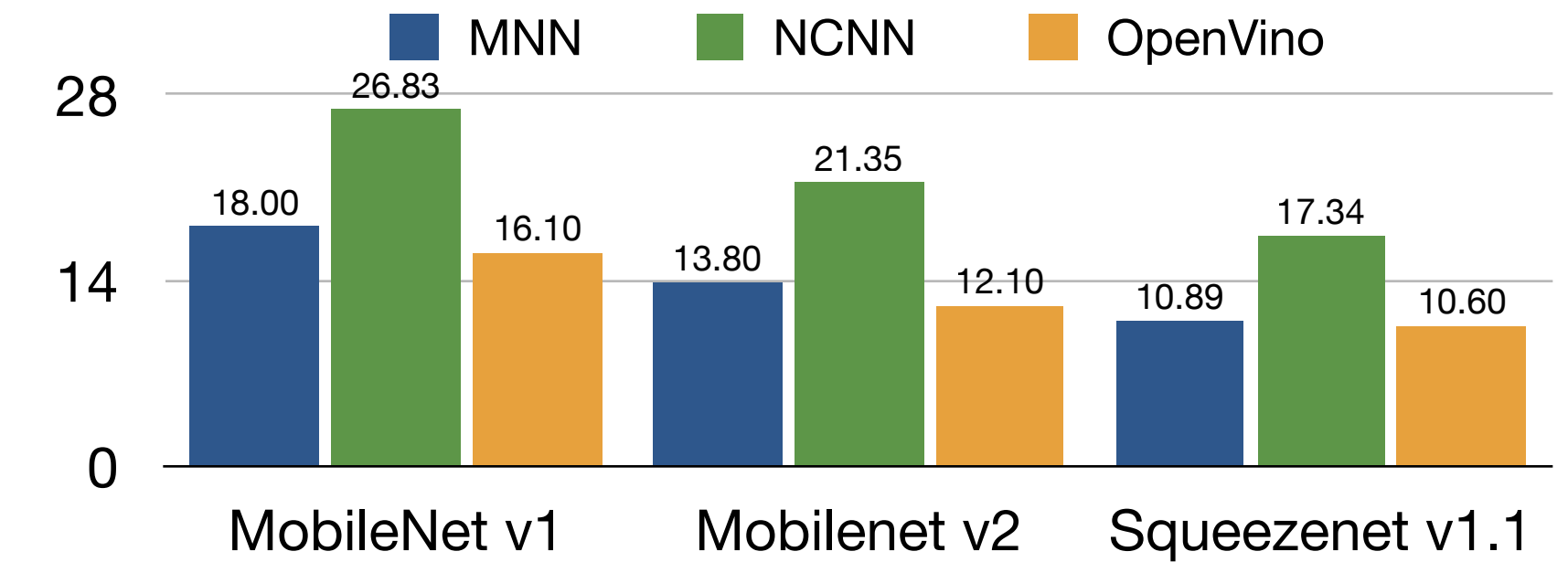
Model: CV / NLP / ASR /...			Framework	Hardware	Device	System
Inception	Mobilenet	Shufflenet	 	 	 	 
Yolo	NanoDet	GAN		 	 	 
RNN	Bert	Transformer				
.....						

MNN - A Universal and Fast Inference Engine

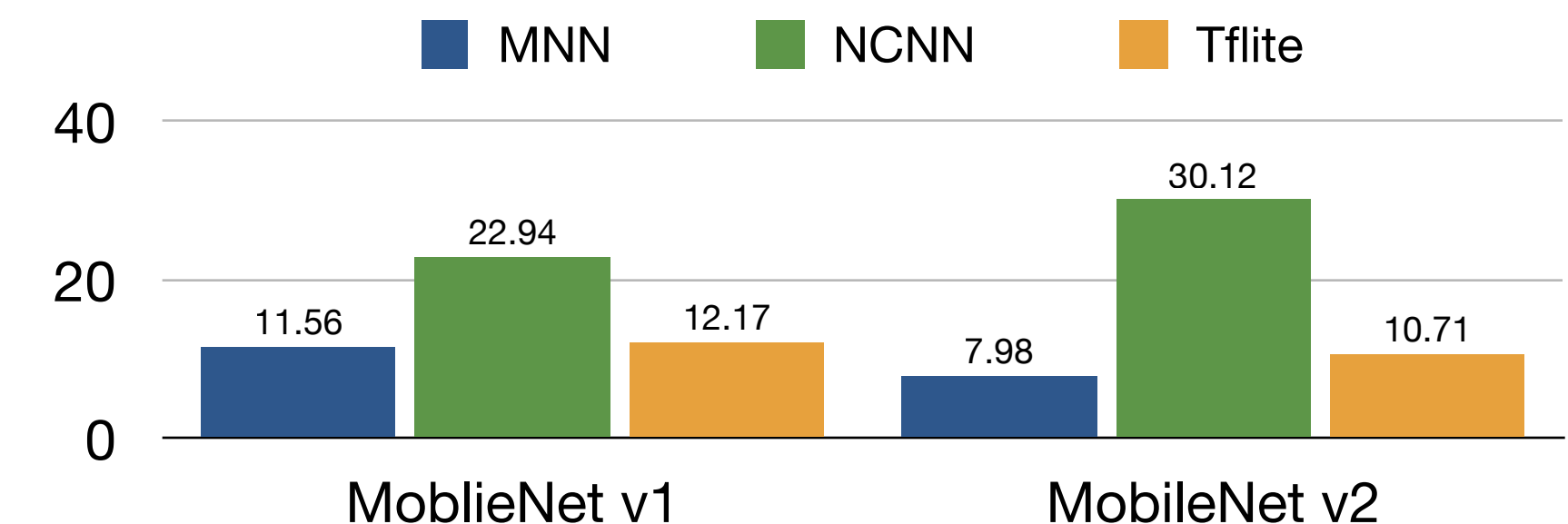
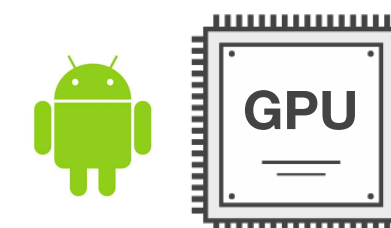
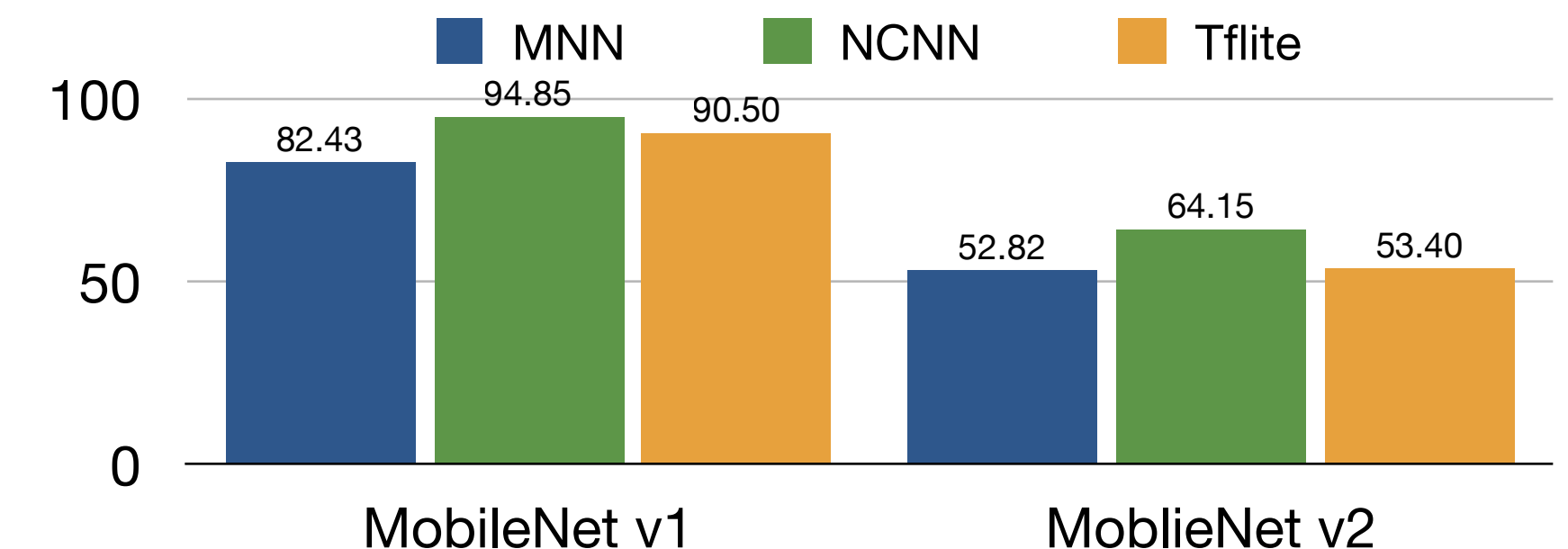
- Method for High Performance
 - Op Fusion
 - NC4HW4 Memory Layout
 - Winograd Convolution / Deconvolution
 - Strassen Matrix Multiply
 - Hand-write Assembly
 - SIMD: NEON / SSE / AVX2 / AVX512
 - CL / GLSL / CUDA
 - Low-precision: FP16 / BF16 / Int8
 - Auto-Tuning



Mac



Mi6



User Case



AR



Live



Image Recognition

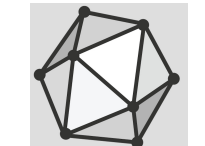


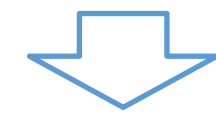
Speech Recognition

Challenge of Universal Inference Engine

Models	Resnet	Inception	VGG	Mobilenet	Shufflenet
	Fastrcnn	Yolo	SSD	NanoDet	GAN
	RNN	LSTM	Bi-LSTM	Bert	Transformer

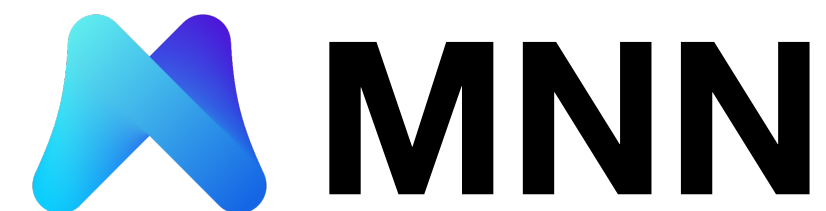
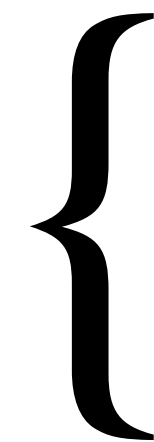
Frameworks

Caffe  Onnx  Tensorflow 

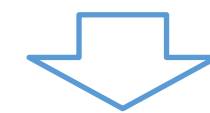


Various operator from different framework

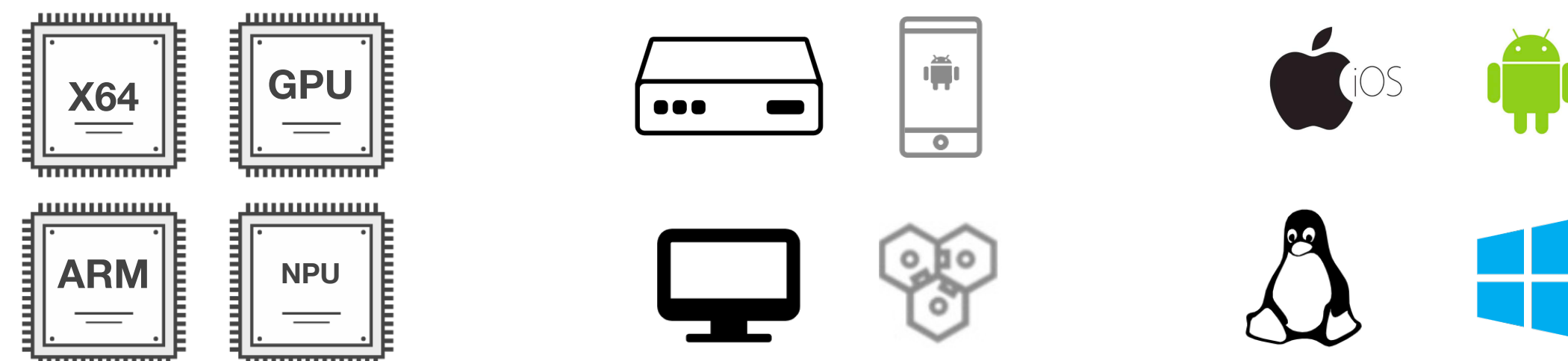
Challenge



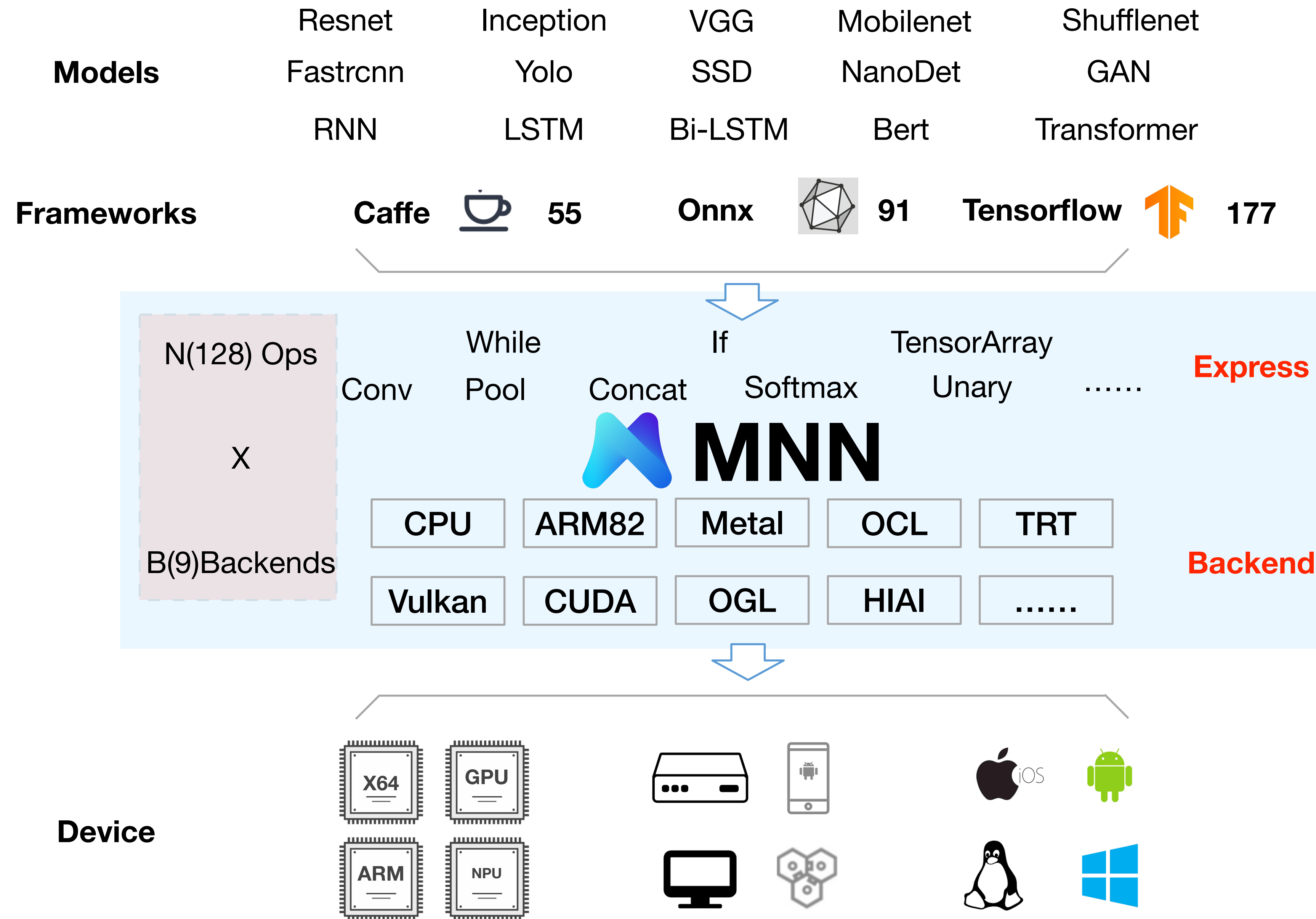
Various Device for operator implement
Different API for using same Hardware



Device



Challenge of Universal Inference Engine

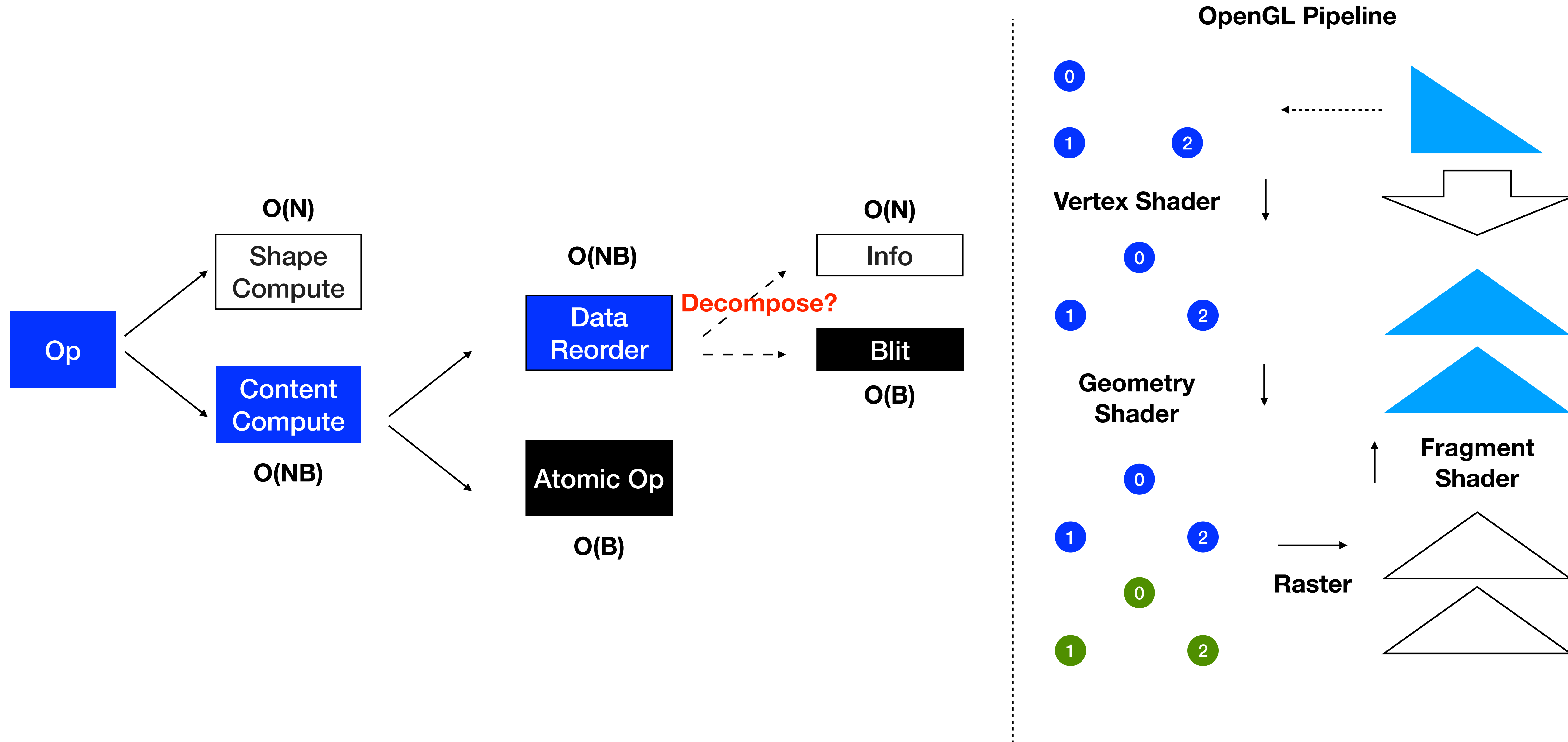


Classification for Operators

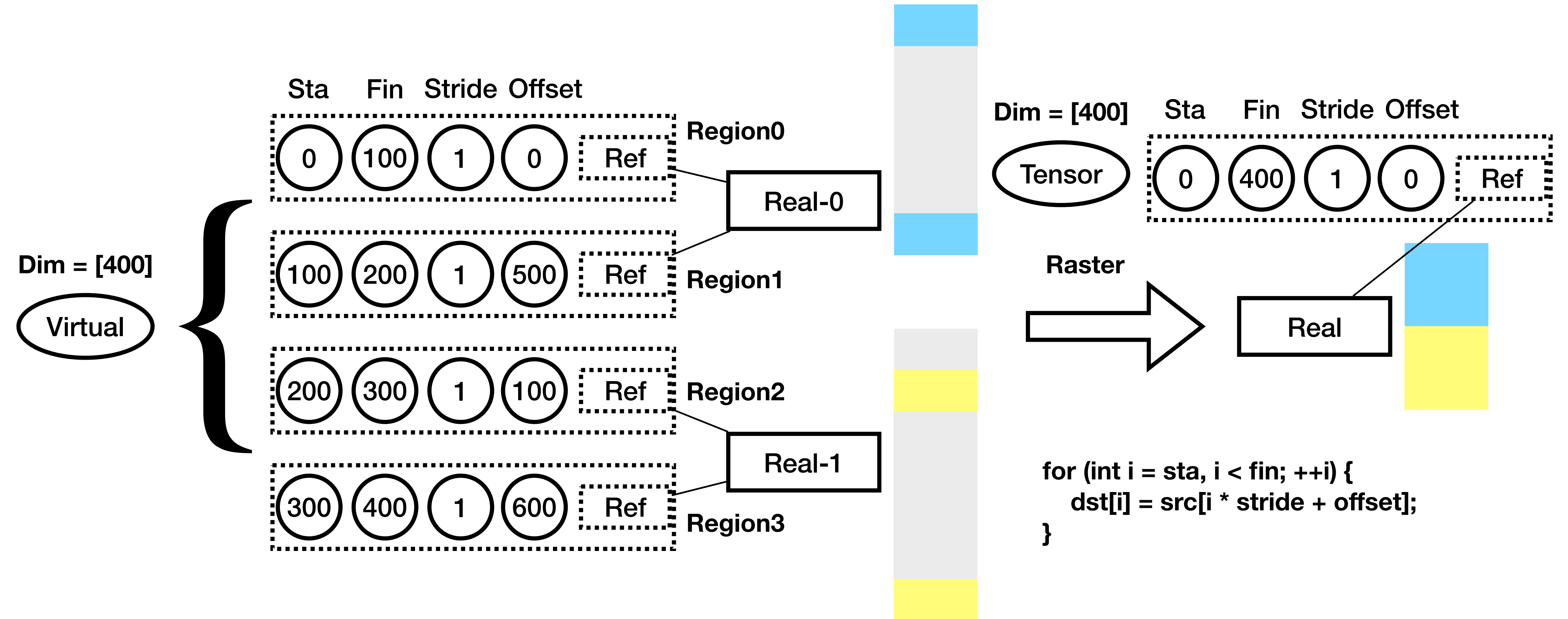


MNN-Express Support	Control flow			
	While		If	
Heavy Work O(NB)	Data Reorder		Fused	
	TensorArrayGather	TensorArraySplit	Resize	Conv3D
	TensorArrayScatter	TensorArrayConcat	RNN	Dilate
	Concat	SpaceToDepth	Conv2D	Pool
	Reshape	BatchToSpace	Softmax	Deconv
	Permute	Crop	NMS	LRN
	Pad	Slice	MatMul	LSTM
				
Easy to Support O(B)	Atomic			
	Unary	Reduce	Binary	

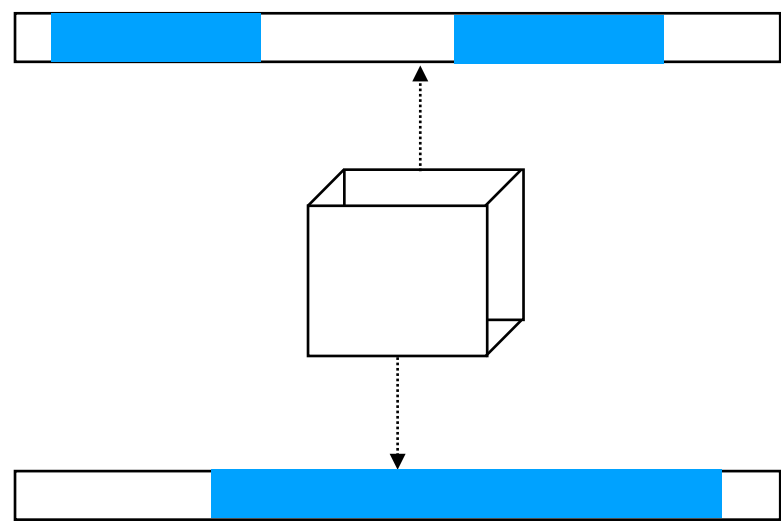
Operator Decomposition



Virtual Tensor and Raster



```
struct View {
    int32_t offset = 0;
    int32_t stride[3] = {1, 1, 1};
};
struct Region {
    View src;
    View dst;
    int32_t size[3] = {1, 1, 1};
    Tensor* origin;
};
```



For x, y, z in range(size):
Dst[x * dst.stride[0]
+ y * dst.stride[1]
+ z * dst.stride[2]
+ dst.offset
] = Src[x * src.stride[0]
+ y * src.stride[1]
+ z * src.stride[2]
+ src.offset]

Data Reorder Op: Compute Region

A: [2, 4]

Transpose
[1, 0]

Virtual B: [4, 2]

Region0:

srcOffset: 0

DstOffset:0

Size: 1, 4, 2

dstStride: 0, 2, 1

srcStride: 0, 1, 4

A: [2, 4]

Slice
[1, 1, 0, 4]

Virtual B: [1, 4]

Region0:

srcOffset: 4

DstOffset:0

Size: 1, 1, 4

dstStride: 0, 4, 1

srcStride: 0, 4, 1

A: [2, 4]

B: [2, 4]

Concat
Axis = 1

Virtual C: [2, 8]

Region0:

srcOffset: 0

DstOffset:0

Size: 1, 2, 4

dstStride: 0, 8, 1

srcStride: 0, 4, 1

Region1:

srcOffset: 0

DstOffset:4

Size: 1, 2, 4

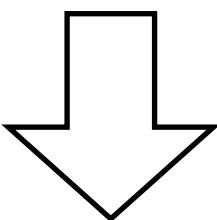
dstStride: 0, 8, 1

srcStride: 0, 4, 1

Fused Op: Decompose

Convolution

$$X(n, c_i, h, w), W(c_o, c_i, k_h, k_w), Y(n, c_o, h, w)$$
$$Y = \text{Conv}(X, W)$$



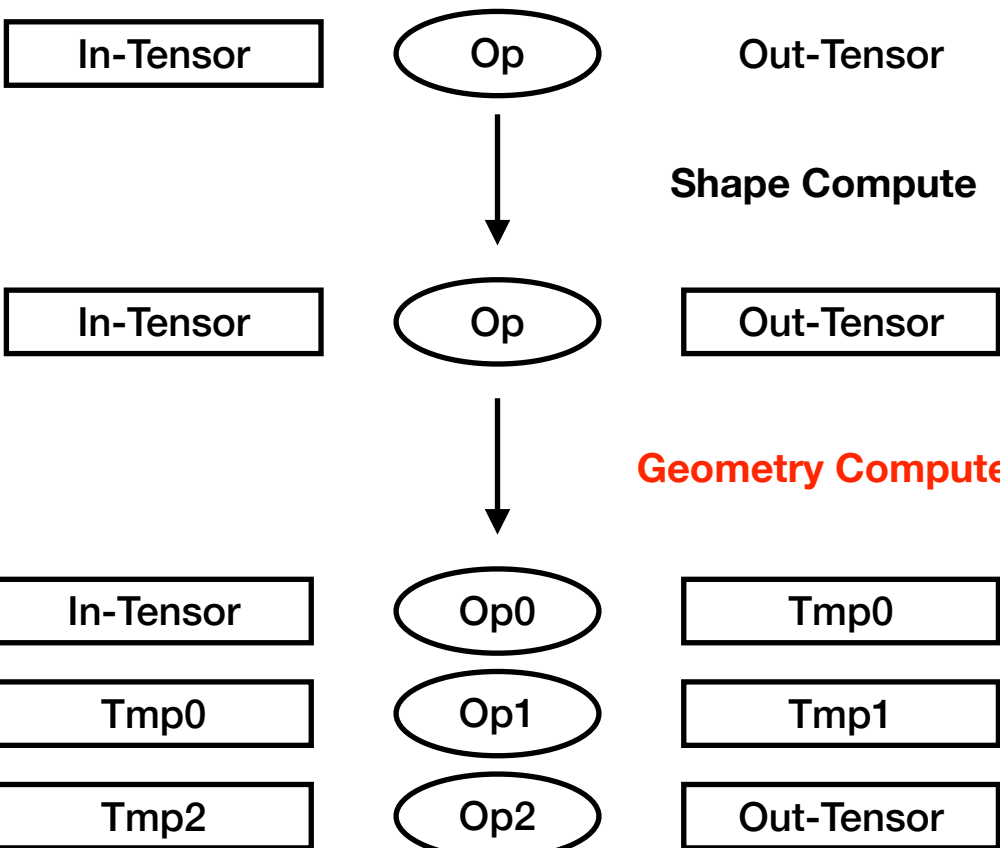
Im2Col + GEMM

$$OX = \text{Reorder}(X)([n, c_i, h, w] \rightarrow [nhwc_o, c_i k_h k_w])$$

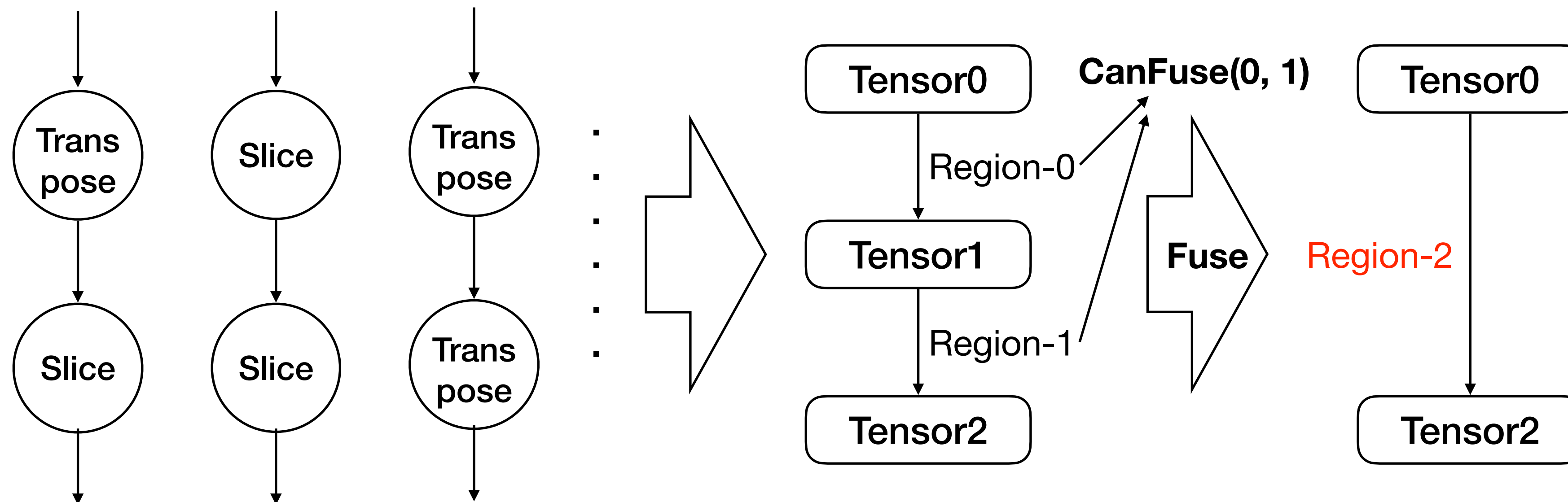
$$OW = \text{Reorder}(W)([c_o, c_i, k_h, k_w] \rightarrow [c_o, c_i k_h k_w])$$

$$OY = \text{MatMul}(OX, OW)$$

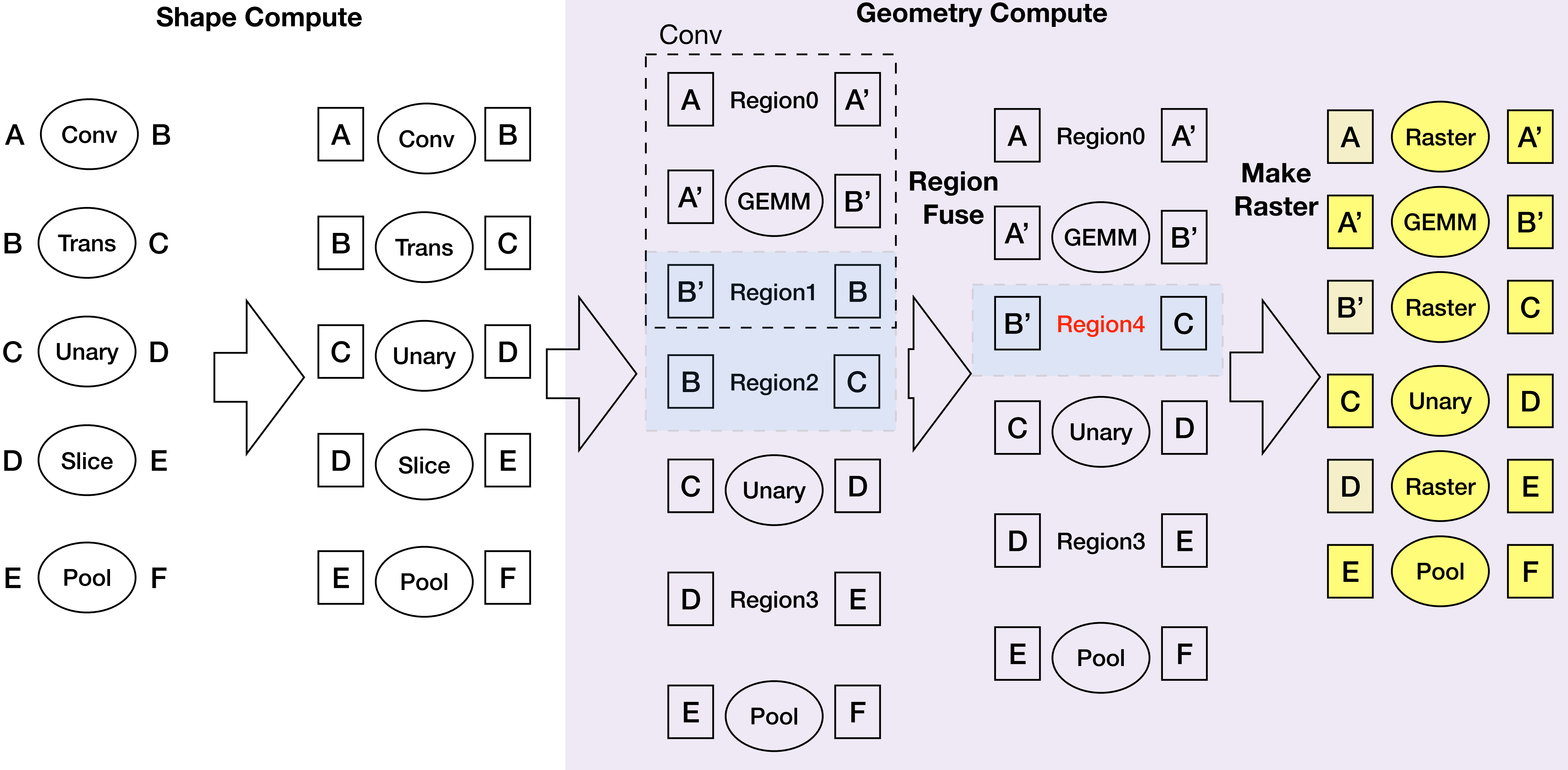
$$Y = \text{IndexMapping}(OY)$$



Region Fuse



Complement in Inference Engine



New Architecture Based on Geometry Compute



Traditional Architecture

$$T = C_{bn}BN$$

CPU

Convolution	Transpose
Deconvolution	Slice
LSTM	Unary
GRU	Reduce
MatMul	Concat
Binary	BatchToSpace

OCL

Convolution	Transpose
Deconvolution	Slice
LSTM	Unary
GRU	Reduce
MatMul	Concat
Binary	BatchToSpace

DSP

Convolution	Transpose
Deconvolution	Slice
LSTM	Unary
GRU	Reduce
MatMul	Concat
Binary	BatchToSpace

CUDA

Convolution	Transpose
Deconvolution	Slice
LSTM	Unary
GRU	Reduce
MatMul	Concat
Binary	BatchToSpace

Geometry Compute Architecture

$$T = C_g(N - N_o) + C_{bn}BN_o$$

Geometry

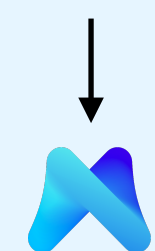
Convolution Deconvolution LSTM GRU Transpose Concat Slice BatchToSpace	CPU	Unary	OCL	Unary
		MatMul		MatMul
		Binary		Binary
		Reduce		Reduce
		Raster		Raster
	DSP	Unary	CUDA	Unary
		MatMul		MatMul
		Binary		Binary
		Reduce		Reduce
		Raster		Raster

Data Reorder Op use Geometry (No Cost)

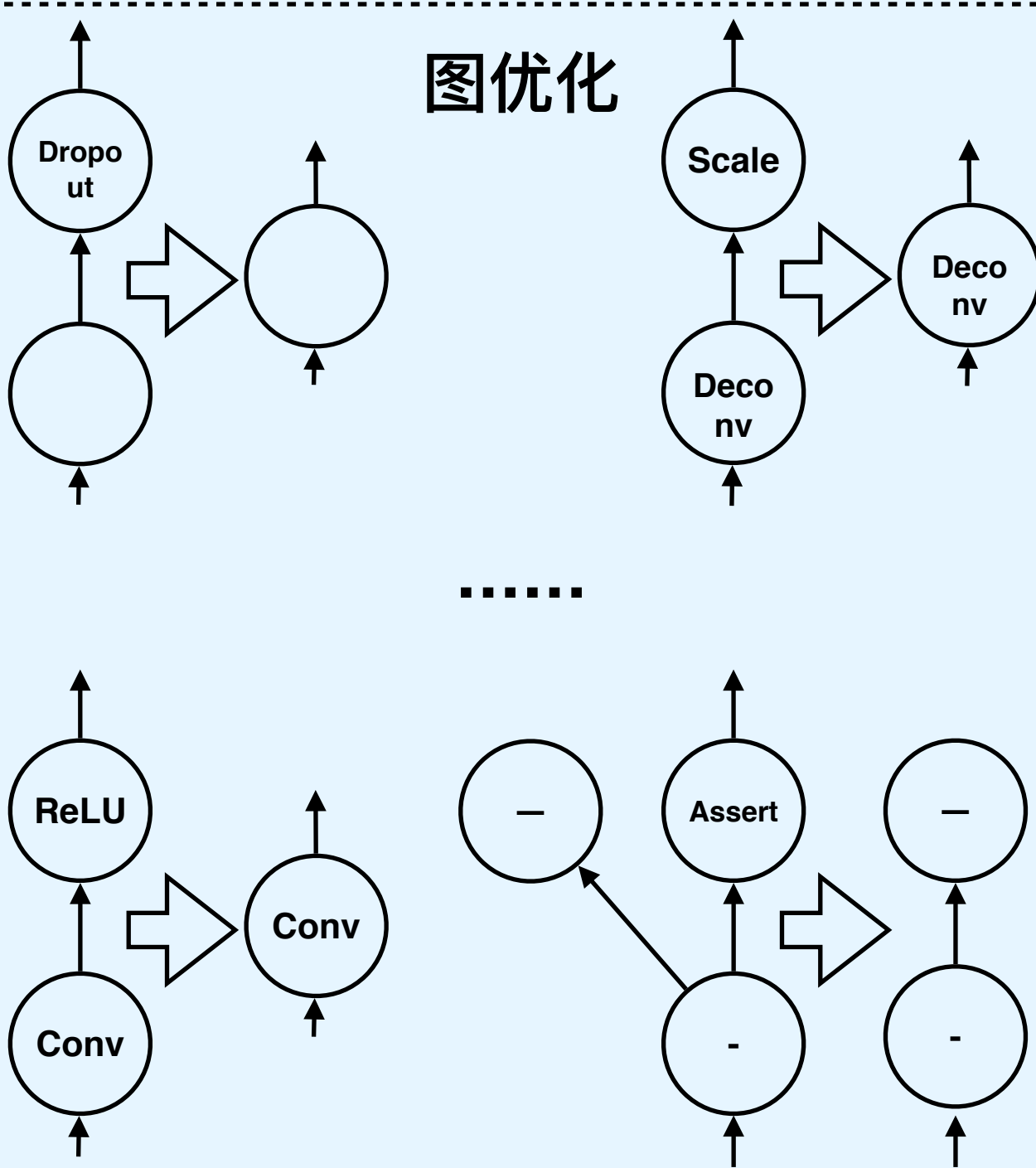
Fused Op decompose (Has Extra Cost)

Core Fused Op hold

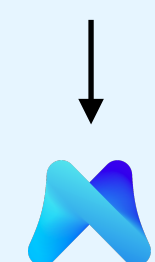
MNN Full Work flow



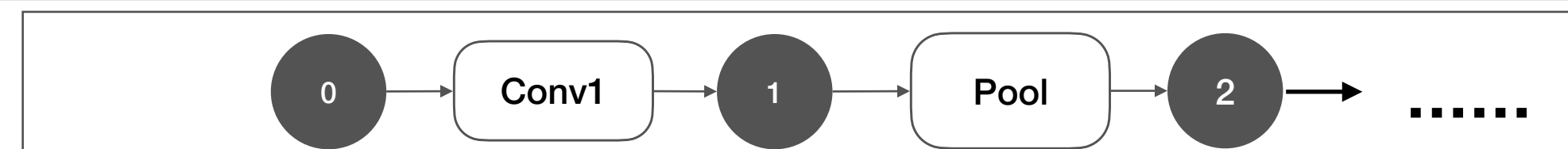
图优化



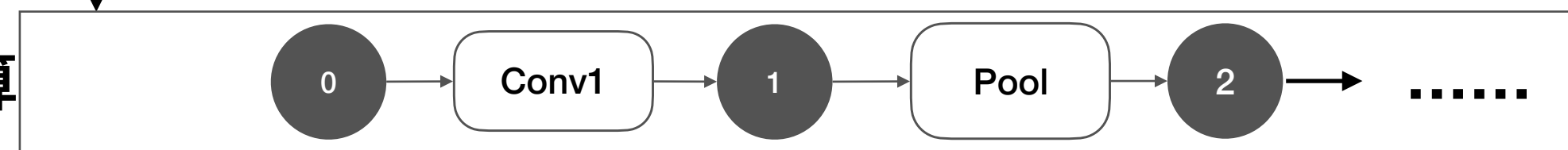
.....



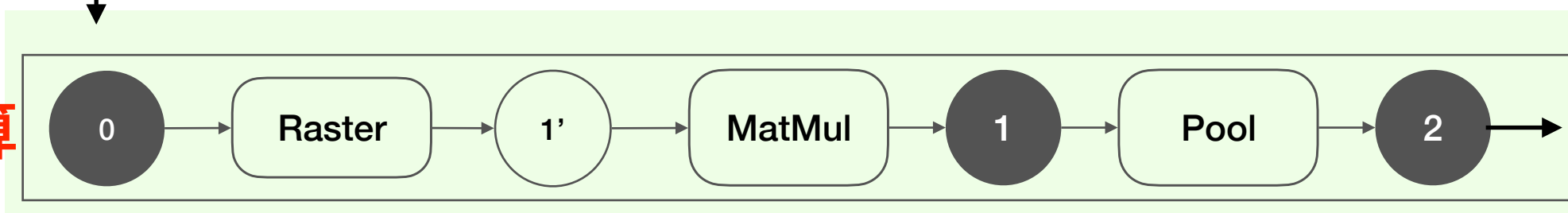
形状计算



用户指定 [1,3,224,224]



几何计算



设备信息

设备	优选后端
Pixels2	Vulkan
Mi6	OpenCL
Mate20	ARM

计算信息

策略搜索

A:算法
B:后端

P:算力
F:计算量

可选计算策略Cost计算

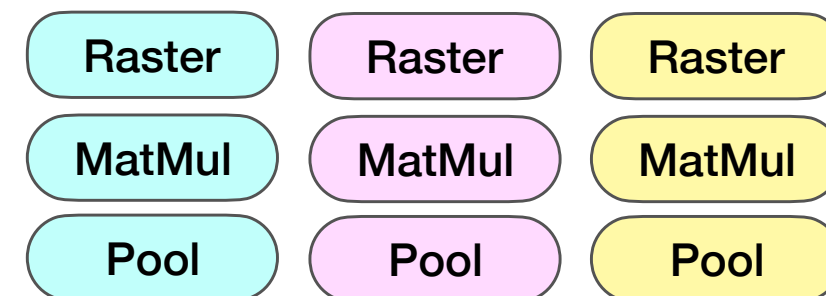
$$C_{B,A} = \frac{F_A}{P_B} + \sigma_A \sigma_B$$

计算步骤
调度损耗

遍历选出Cost最低计算策略(算法+后端)

预推理

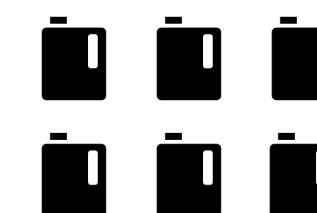
计算策略



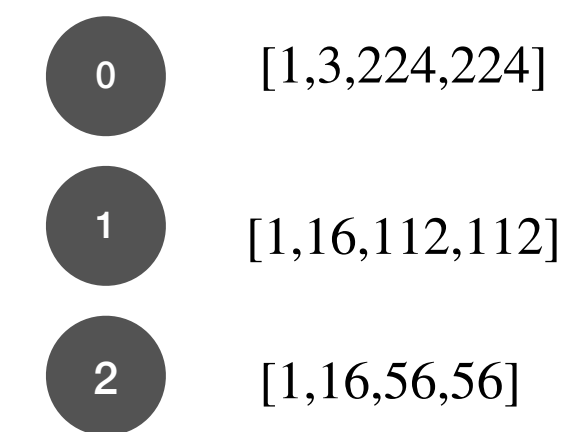
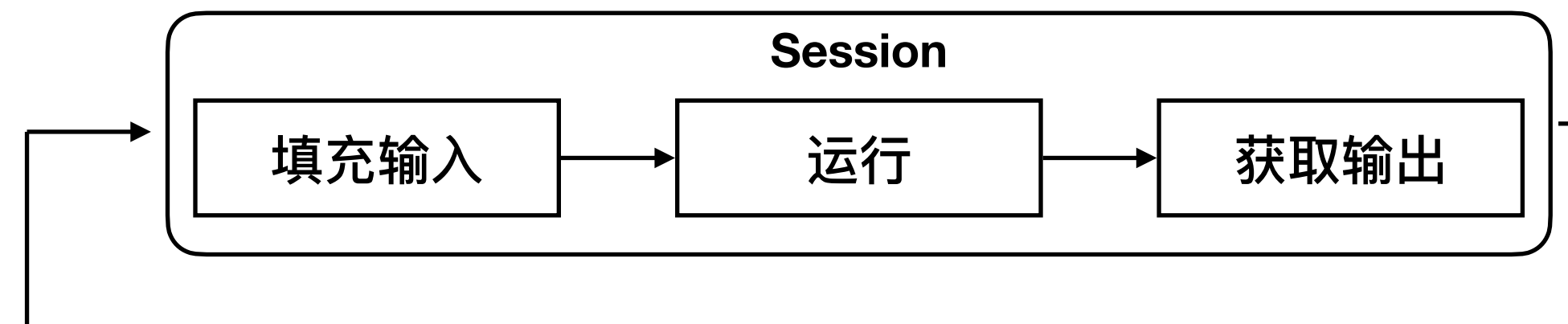
系统资源

(内存/纹理ID/命令缓冲区)

申请资源



推理



形状信息

坐标映射

$$(i', j', k') = f(i, j, k)$$
$$(i', j', k', l') = g(i, j, k, l)$$
$$(i') = h(i, j)$$

区域信息

区域计算

$$R_0$$
$$R_0, R_1, \dots, R_n$$
$$R_0, R_1, \dots, R_m$$

Convolution	Transpose	算子拆解	Unary
Deconvolution	Slice		Reduce
LSTM	Unary		MatMul
GRU	Reduce		Binary
MatMul	Concat		Raster
Binary	BatchToSpace		

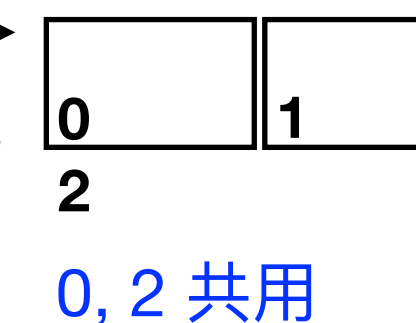
资源分配优化

申请顺序

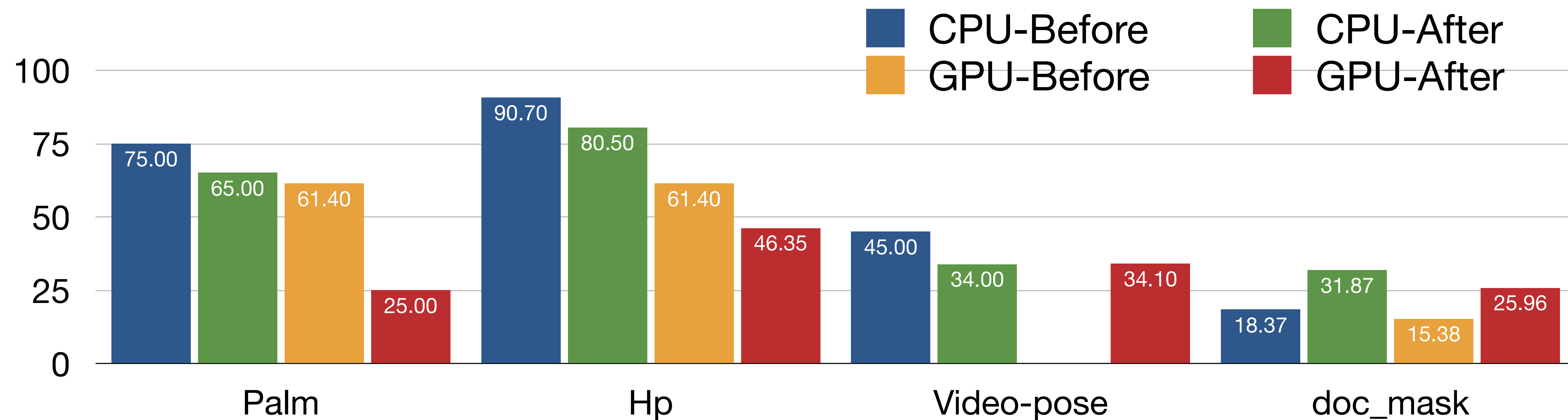
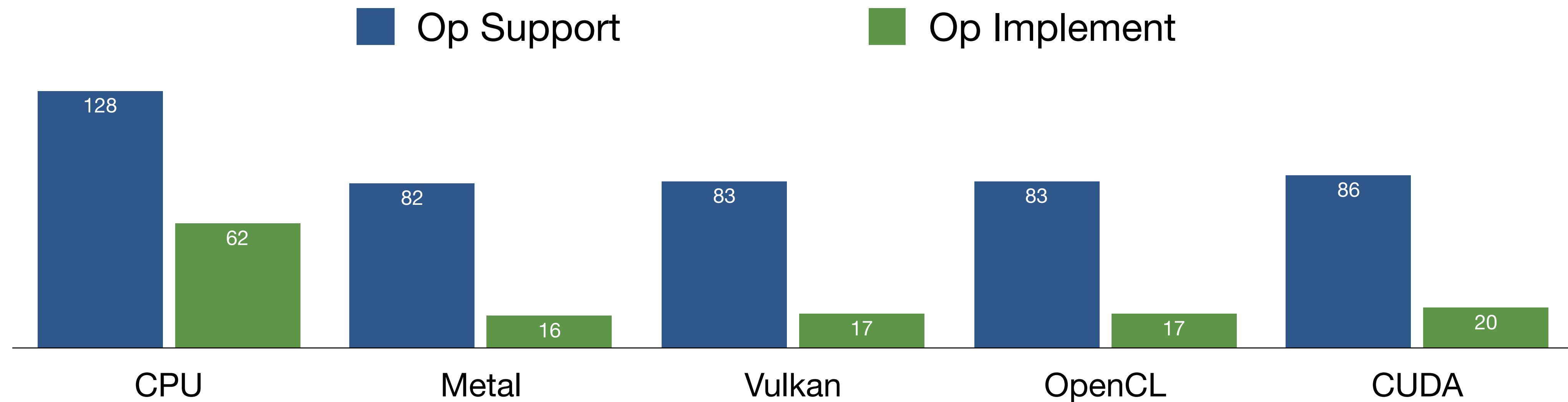
Alloc 0
Alloc 1
Free 0
Alloc 2
Free 1
Free 2

计算最小分配

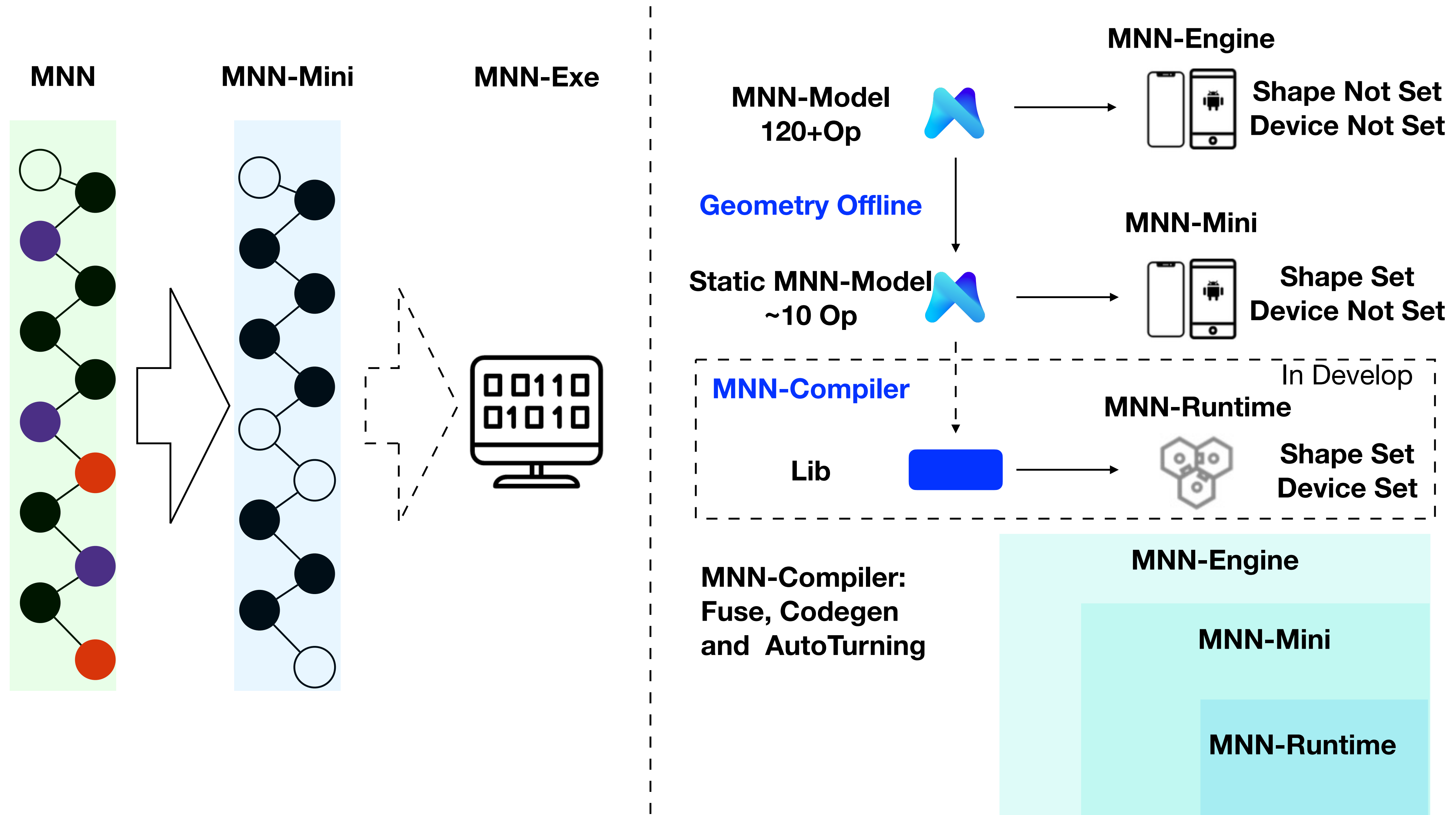
实际分配



Result of Geometry Compute Refract



Geometry for Multi-Level Deploy



- <https://github.com/alibaba/MNN>

- Geometry Compute
 - Speed: Cost Less Decomposition for Fused Op
 - Post-treat/Pre-treat: GIR for more CV / Data Pretreat / Post-treat
 - Train: Grad Based on Geometry Compute
- MNN
 - More Frontend Support, such as Torchscript
 - More Backends for Hardware Support: NPU / DSP
 - Model Compress Algorithm Combined with Compiler Technology



THANKS