

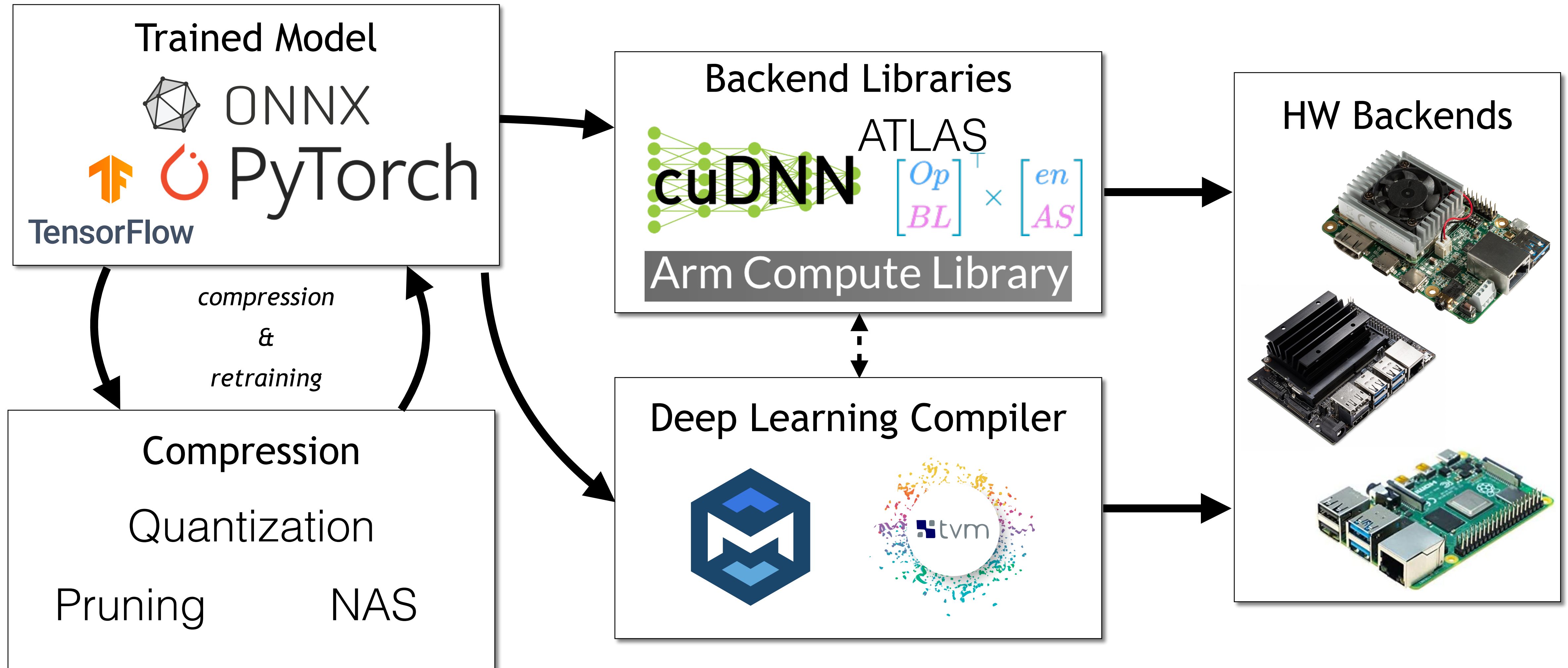


UNIVERSITÄT
HEIDELBERG
ZUKUNFT
SEIT 1386

HARDWARE-AWARE DEPLOYMENT OF (PROBABILISTIC) NEURAL ARCHITECTURES

Bernhard Klein
Computing Systems Group
Institute of Computer Engineering
Heidelberg University

DEPLOYMENT CHAIN: DNNs & EMBEDDED SYSTEMS



GALEN

**ARCHITECTURE-SPECIFIC AUTOMATIC
COMPRESSION**

COMPRESSION POLICY

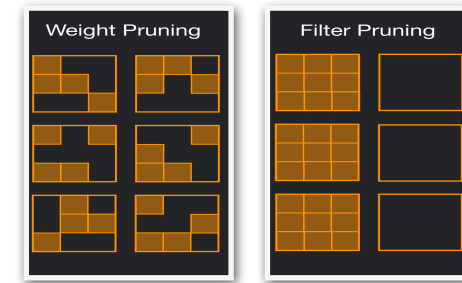


Quantization

Low bit width data types

Pruning

Remove structures



Layer-specific

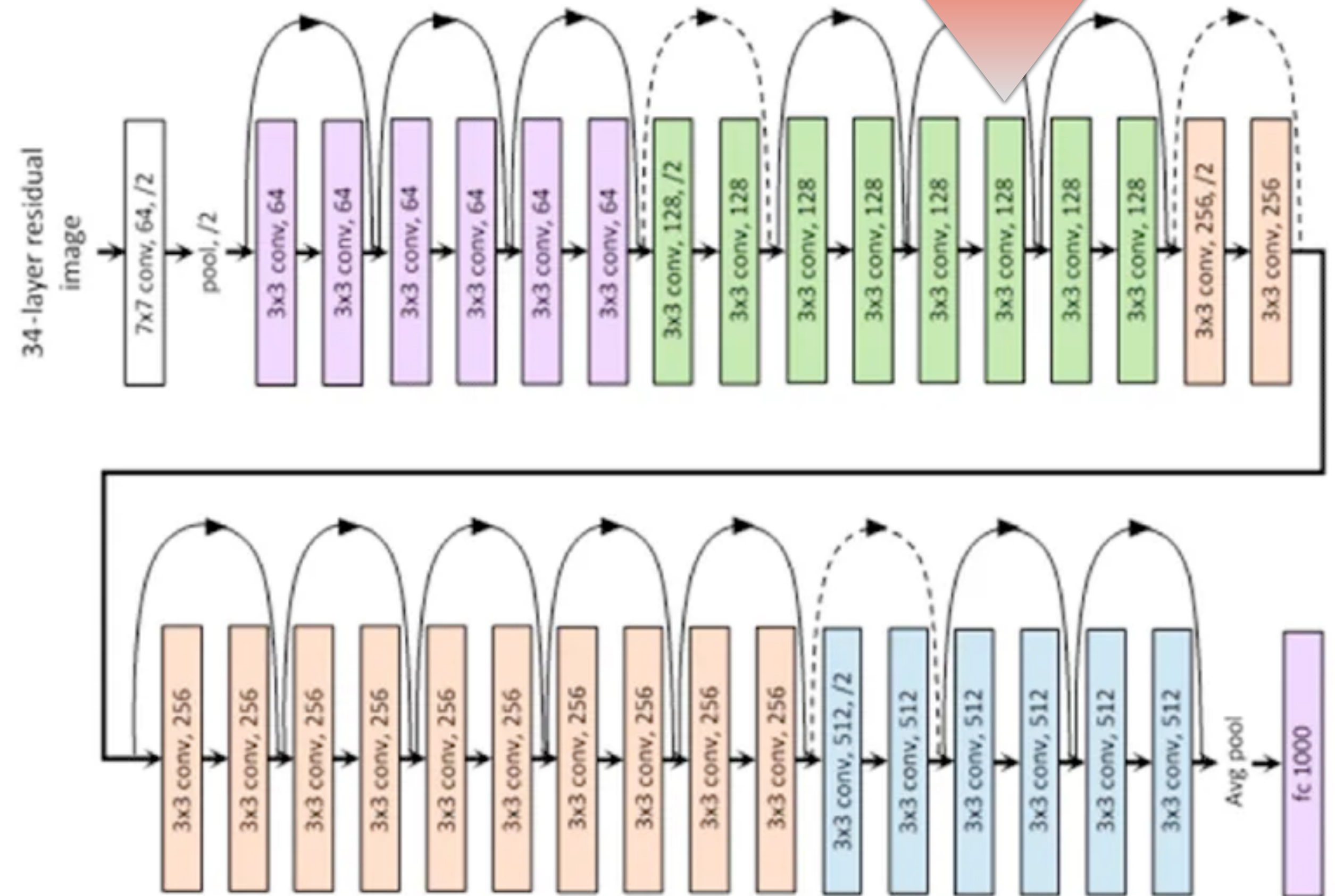
Sparsity

Bit width (activations & weights)

Not all combinations are possible

Affects hardware resources & model accuracy

Model & hardware expertise



COMPRESSION POLICY



Quantization

Low bit width data types

Pruning

Remove structures

Layer-specific

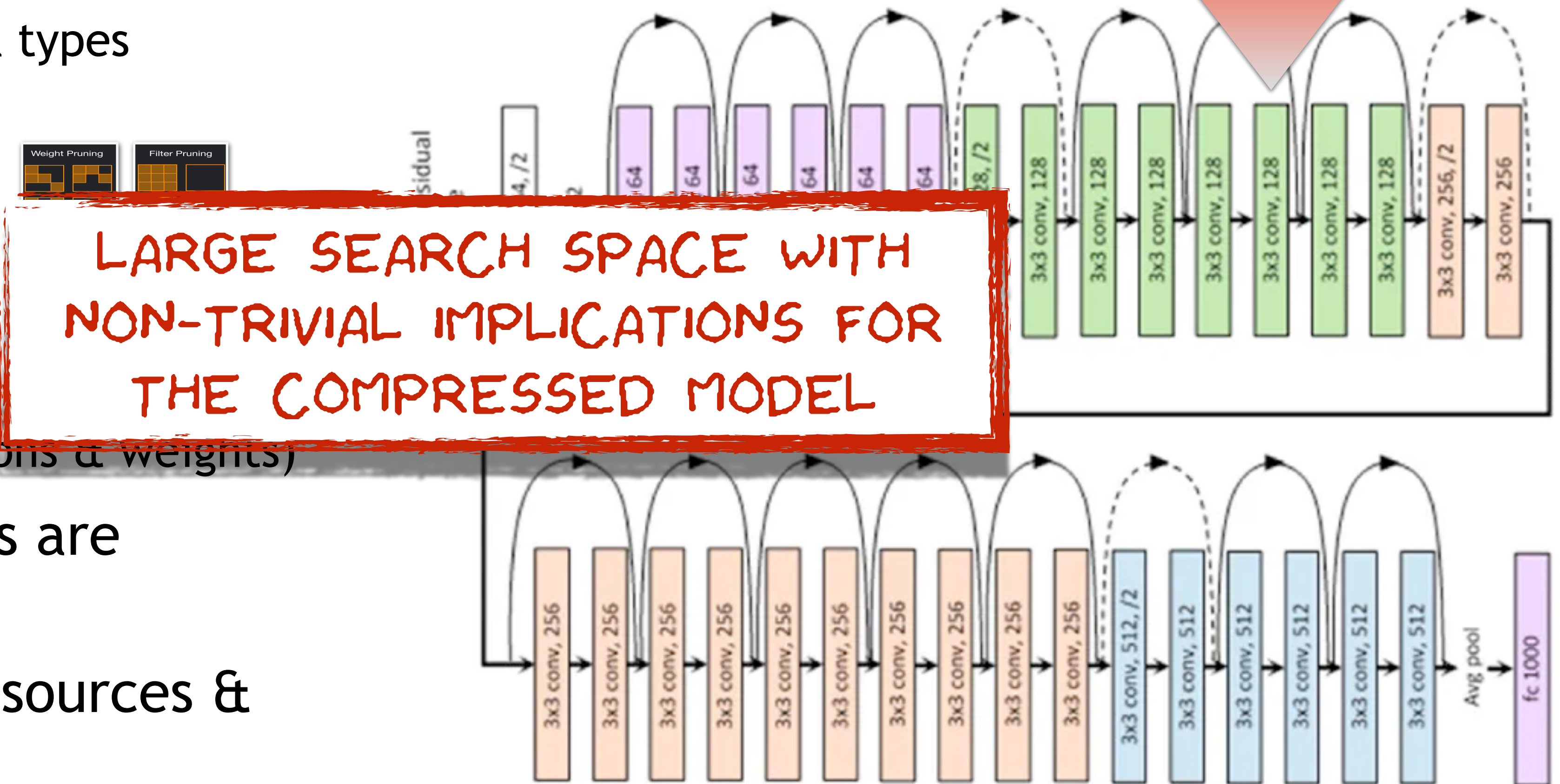
Sparsity

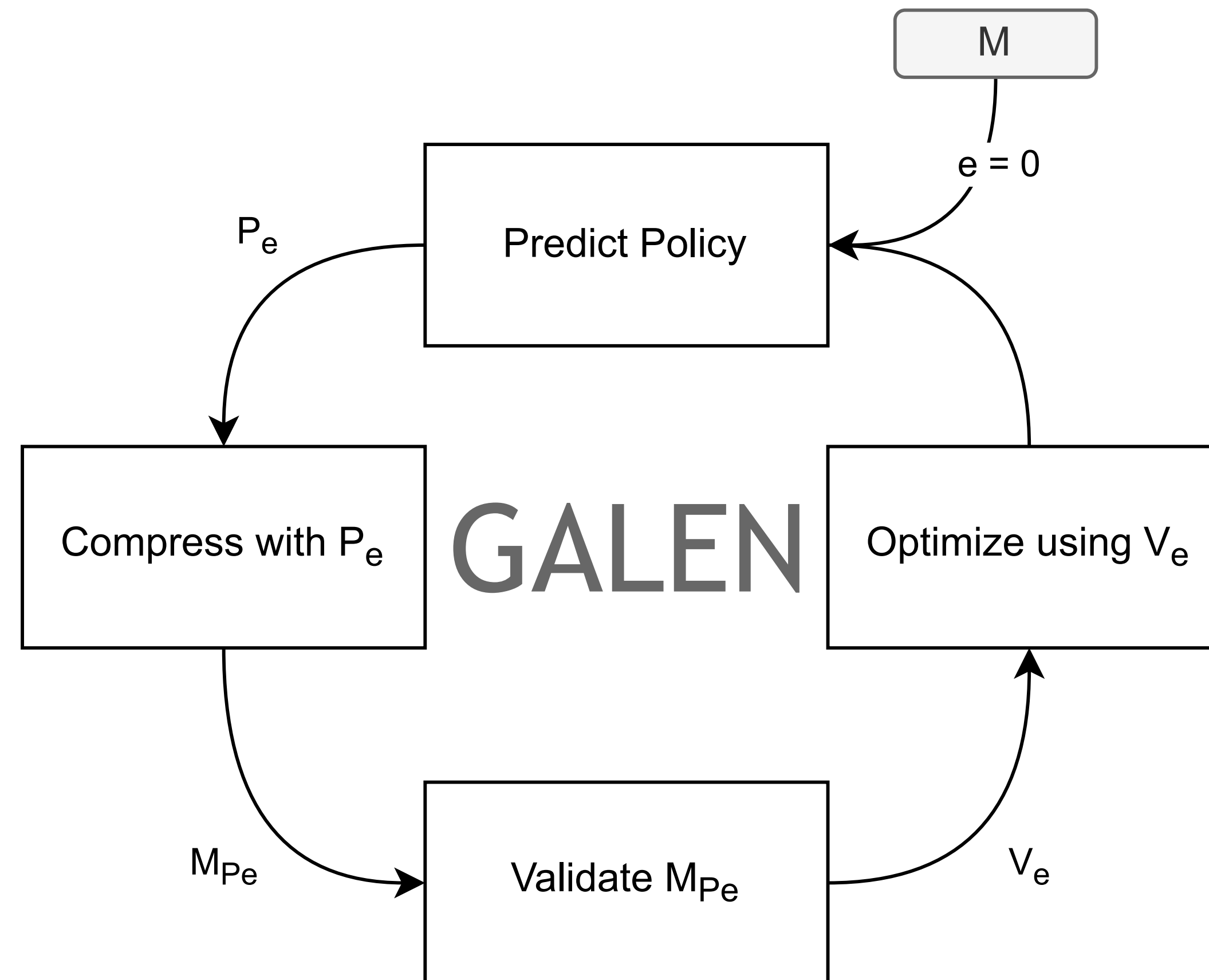
Bit width (activations & weights)

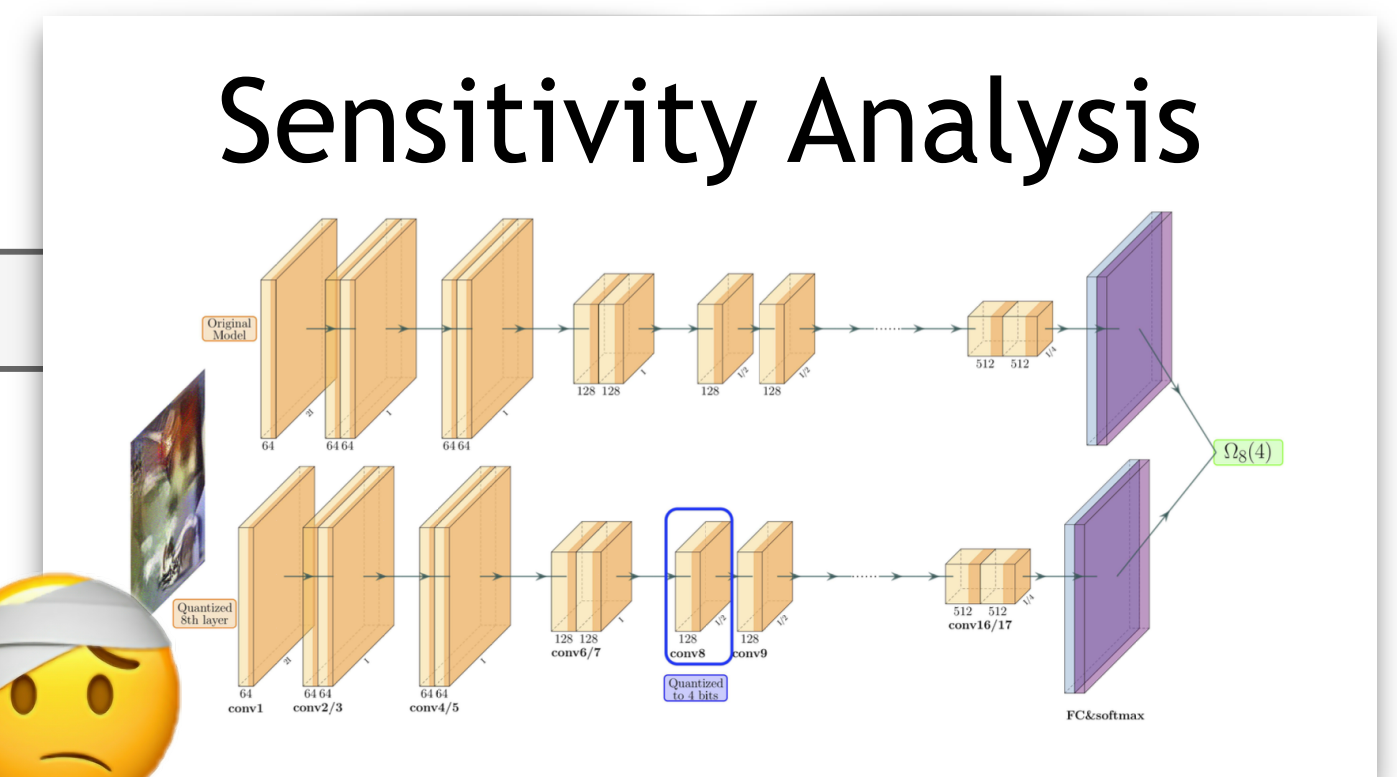
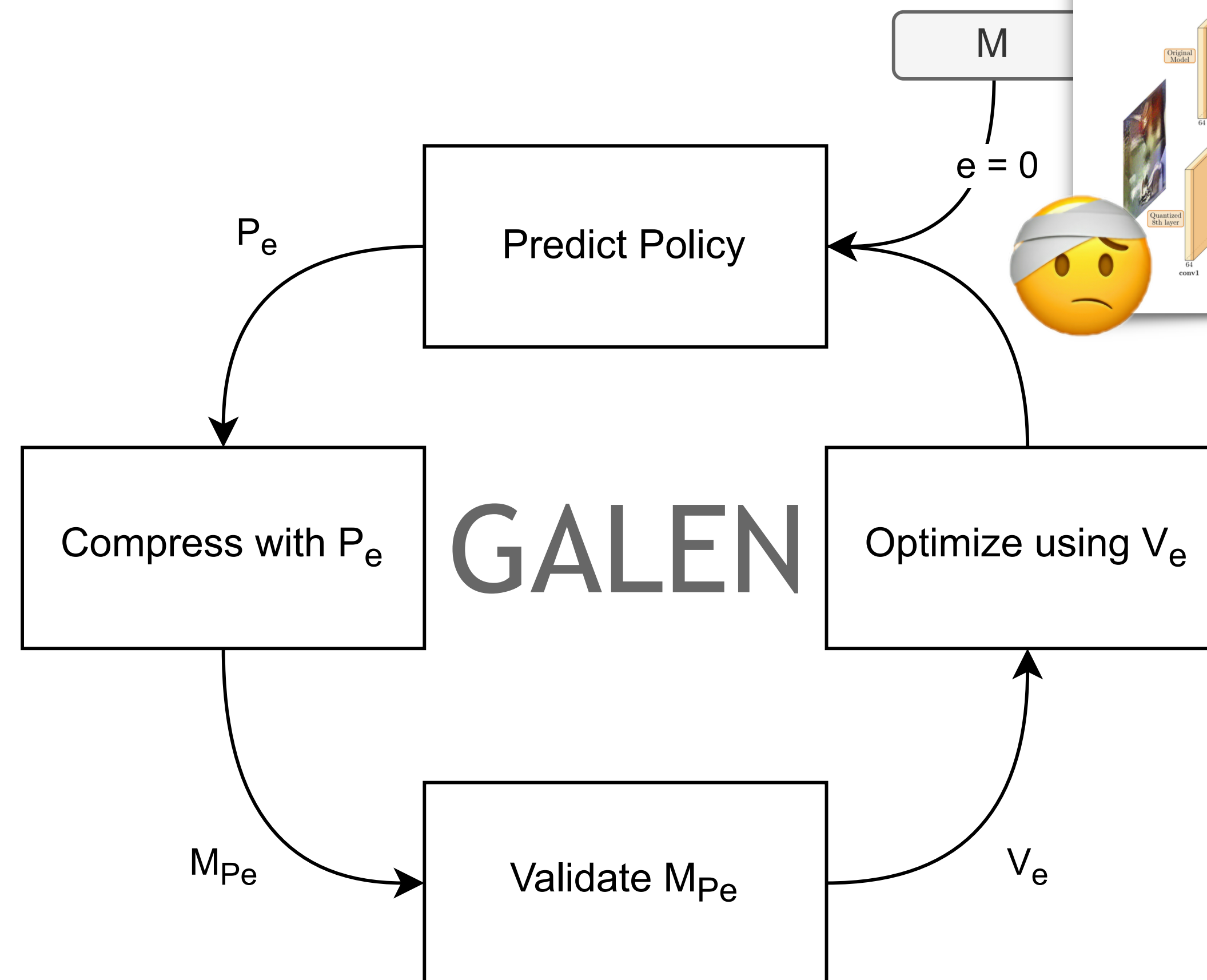
Not all combinations are possible

Affects hardware resources & model accuracy

Model & hardware expertise



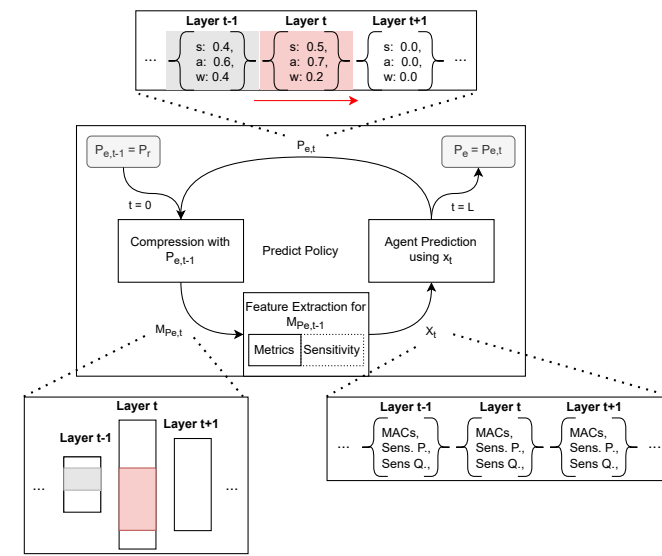




Predict Comp. Policies

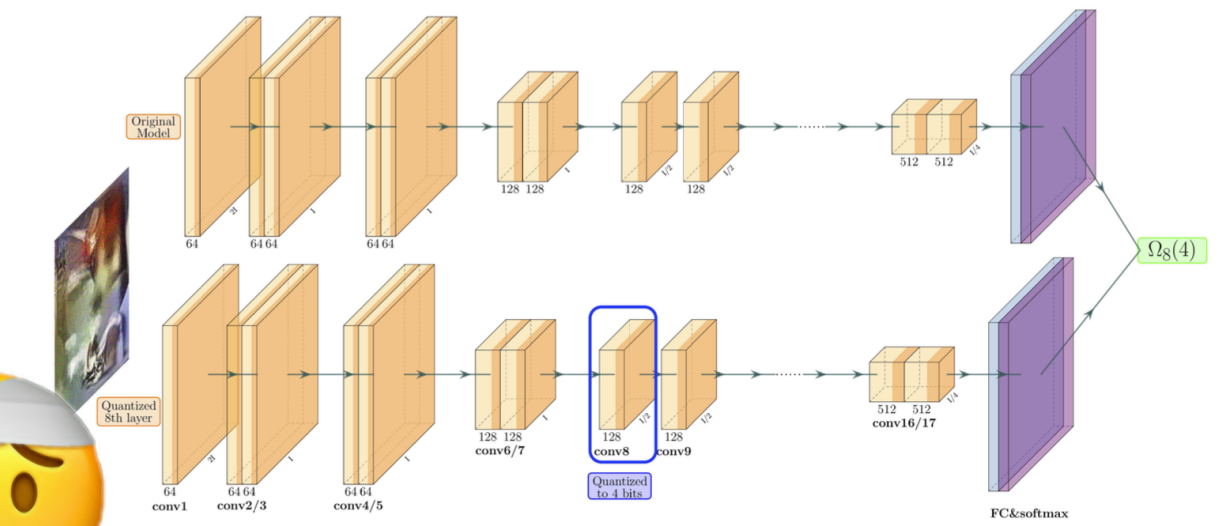
Iterative

Guided
Sensitivity
MACs, BOPs



Predict Policy

Sensitivity Analysis



M

$e = 0$



Compress with P_e

GALEN

Optimize using V_e

M_{P_e}

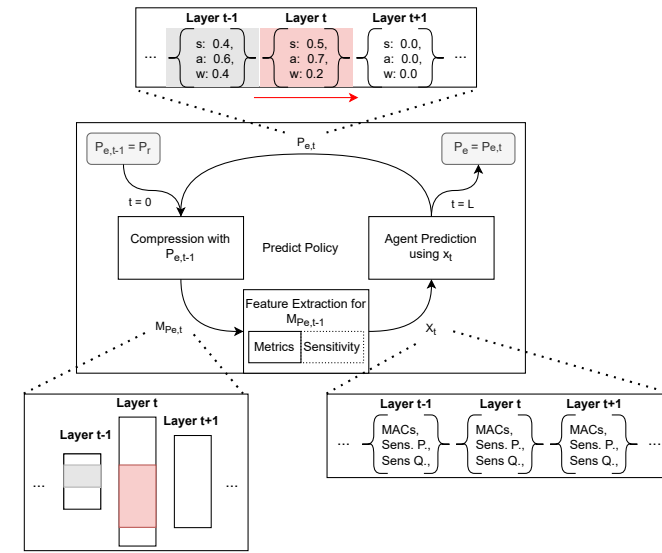
Validate M_{P_e}

V_e

Predict Comp. Policies

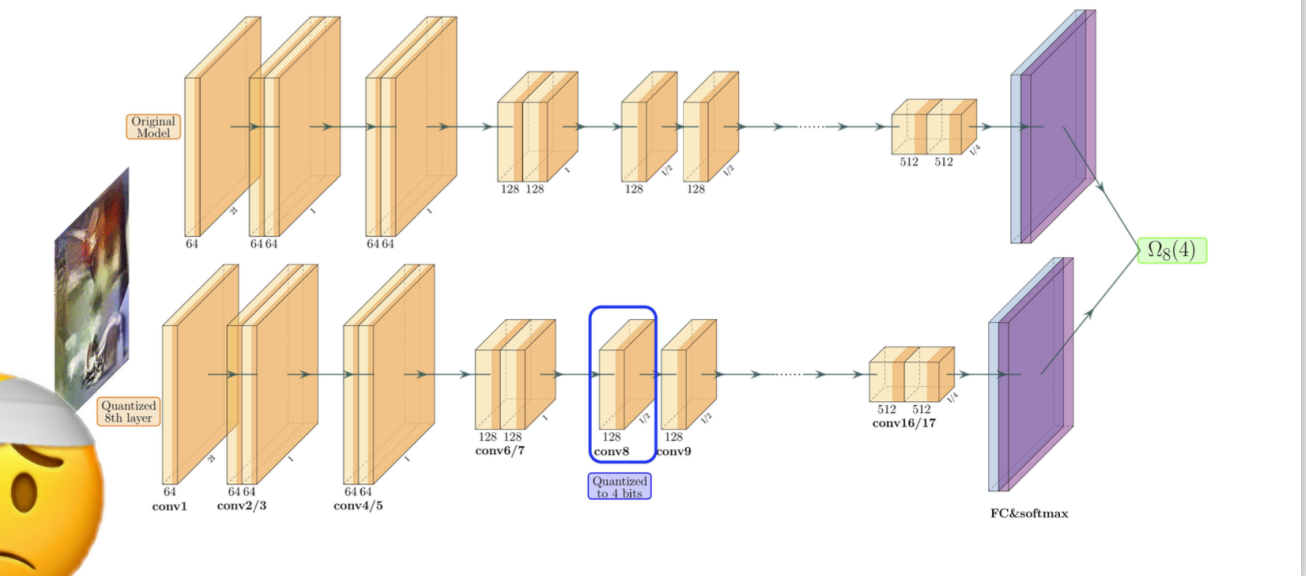
Iterative

Guided
Sensitivity
MACs,BOPs



Predict Policy

Sensitivity Analysis



M

$e = 0$



Pruning

Sparsity $\rightarrow$ Channel Count
Select channels with ℓ_1 strategy

Track dependencies between
layers with Torch-Pruning*

Compress with P_e

GALEN

Optimize using V_e

M_{Pe}

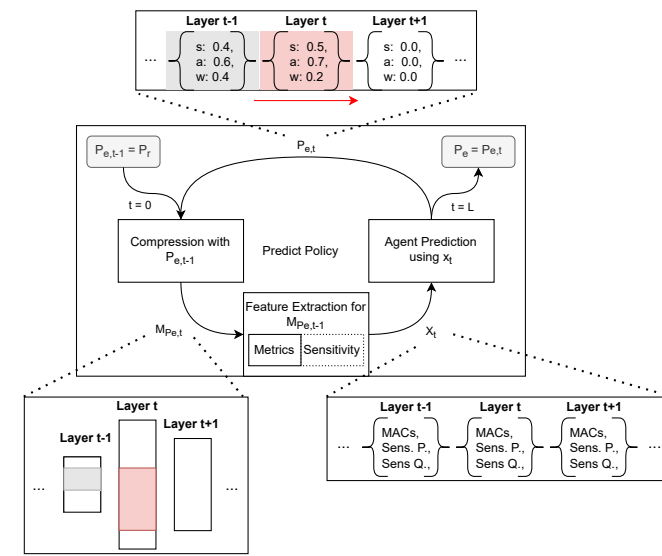
Validate M_{Pe}

V_e

Predict Comp. Policies

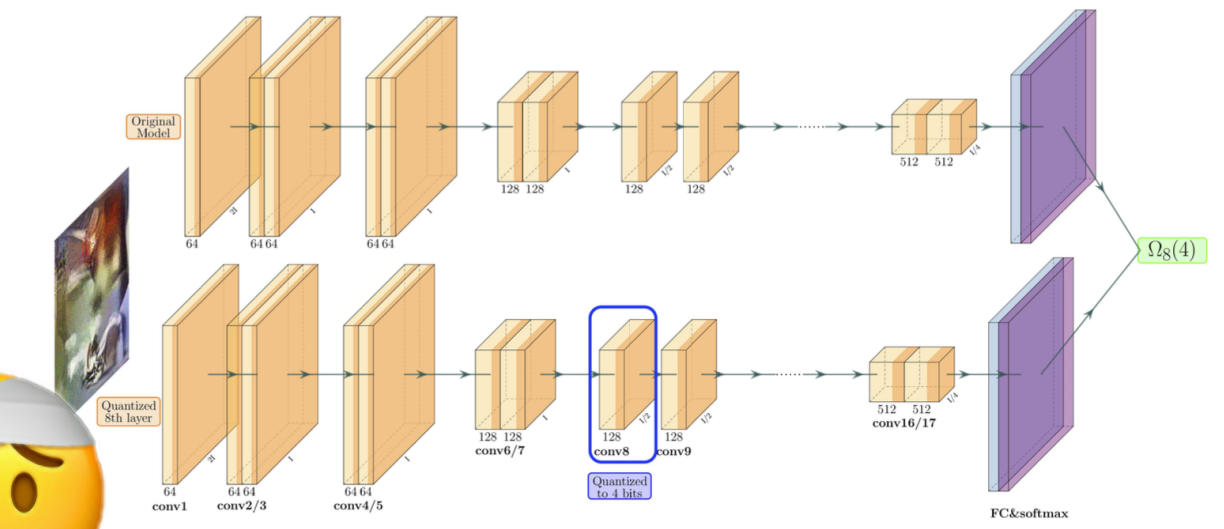
Iterative

Guided
Sensitivity
MACs,BOPs



Predict Policy

Sensitivity Analysis



M

$e = 0$



Pruning

Sparsity $\rightarrow$ Channel Count
Select channels with ℓ_1 strategy

Track dependencies between
layers with Torch-Pruning*

Compress with P_e

GALEN

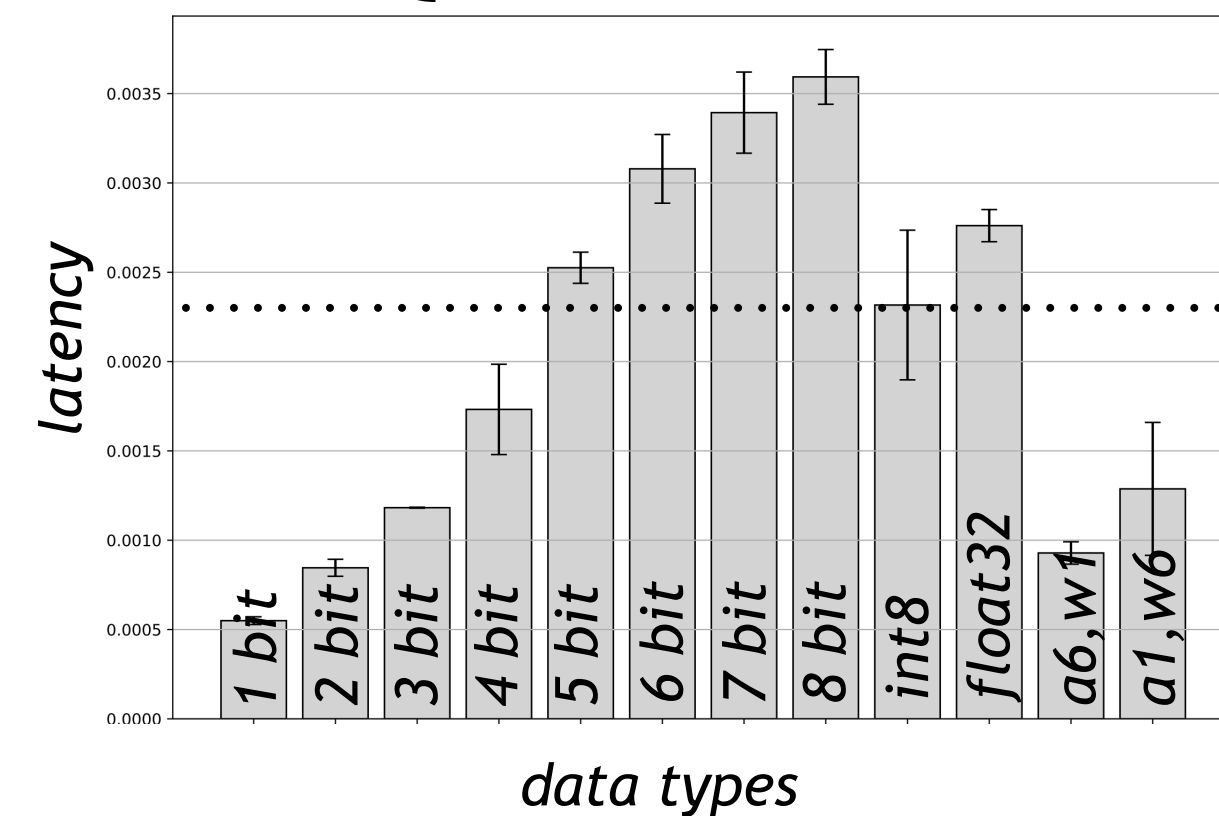
Optimize using V_e

Validate M_{P_e}

M_{P_e}

V_e

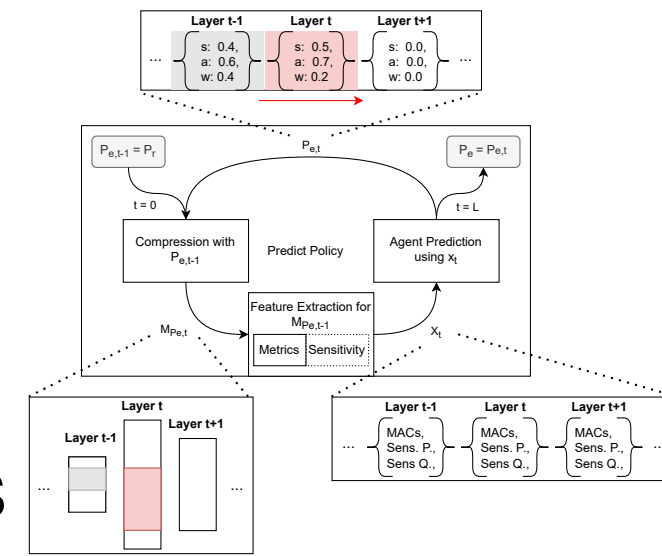
Quantization



Predict Comp. Policies

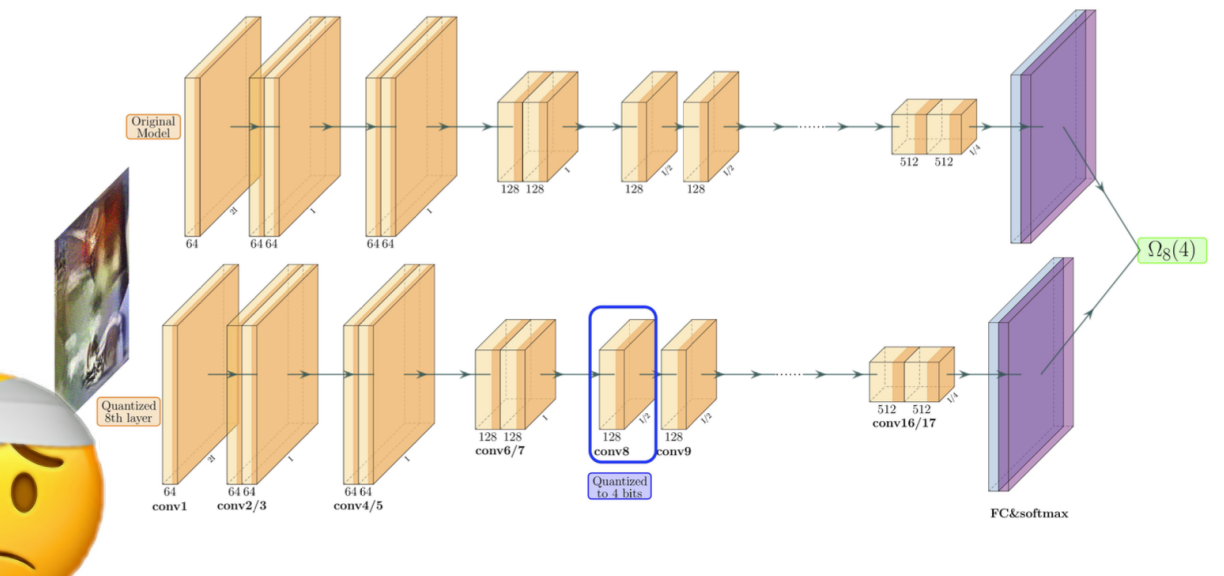
Iterative

Guided
Sensitivity
MACs,BOPs



Predict Policy

Sensitivity Analysis



M

$e = 0$



Pruning

Sparsity $\rightarrow$ Channel Count
Select channels with ℓ_1 strategy

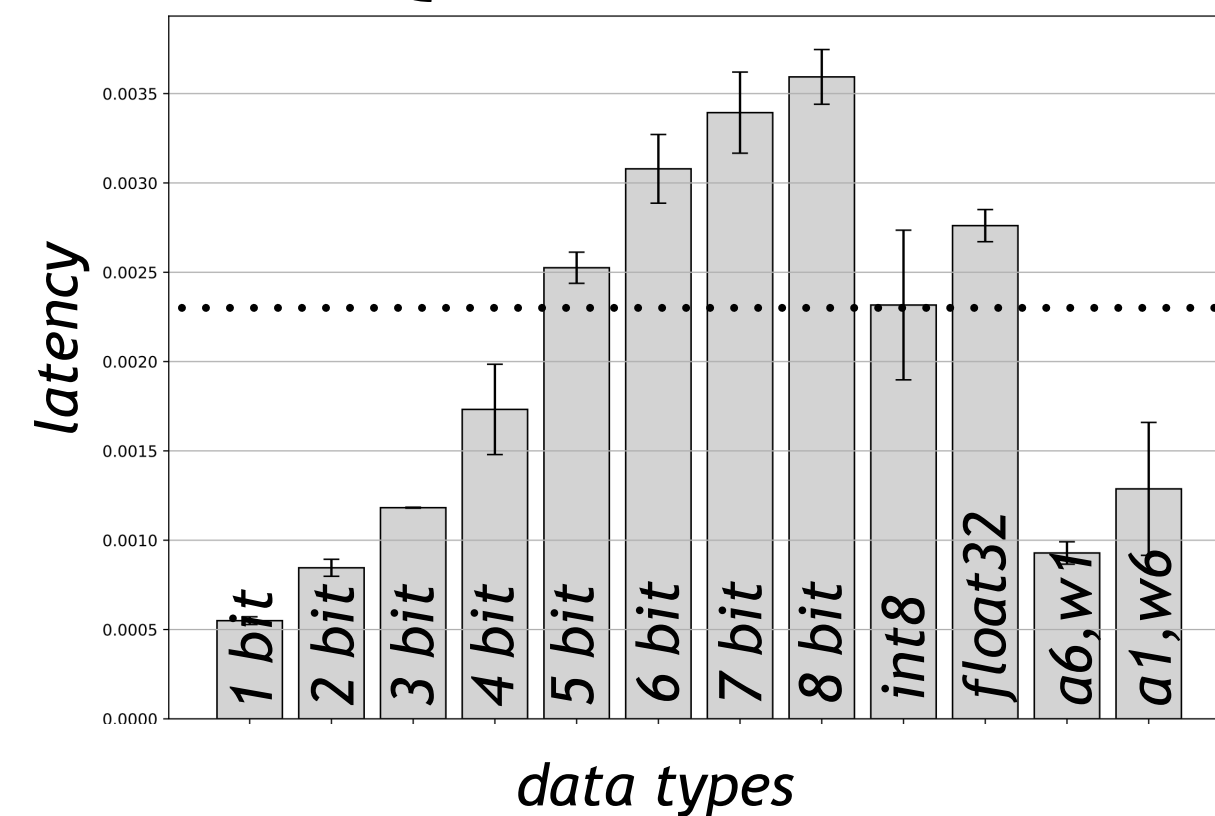
Track dependencies between
layers with Torch-Pruning*

GALEN

Compress with P_e

Optimize using V_e

Quantization



M_{Pe}

Validate M_{Pe}

V_e



Evaluation

Accuracy
(1 ep. retraining)

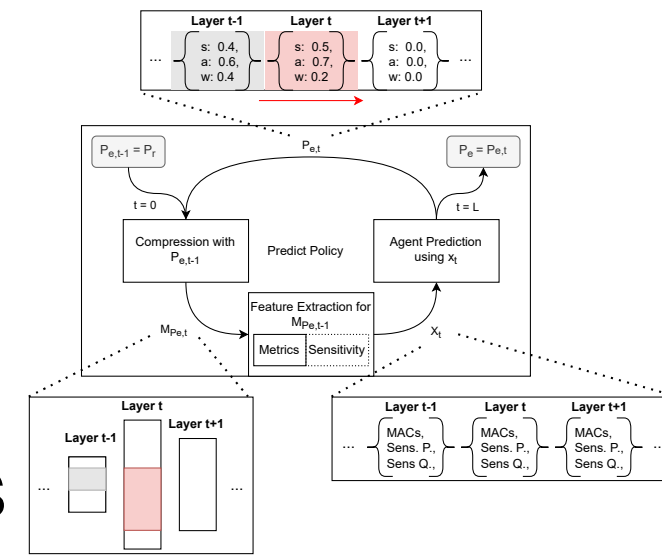
Latency Benchmark
(quantized operators)



Predict Comp. Policies

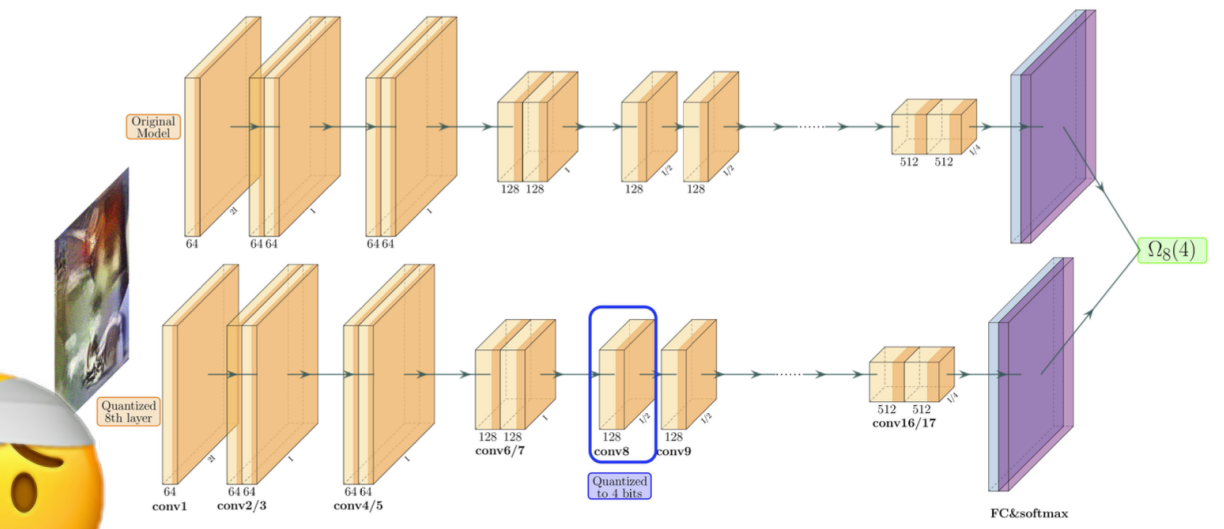
Iterative

Guided
Sensitivity
MACs,BOPs



Predict Policy

Sensitivity Analysis



M

e = 0



Pruning

Sparsity → Channel Count
Select channels with ℓ_1 strategy

Track dependencies between
layers with Torch-Pruning*

Compress with P_e

GALEN

Optimize using V_e

Absolute Reward Function

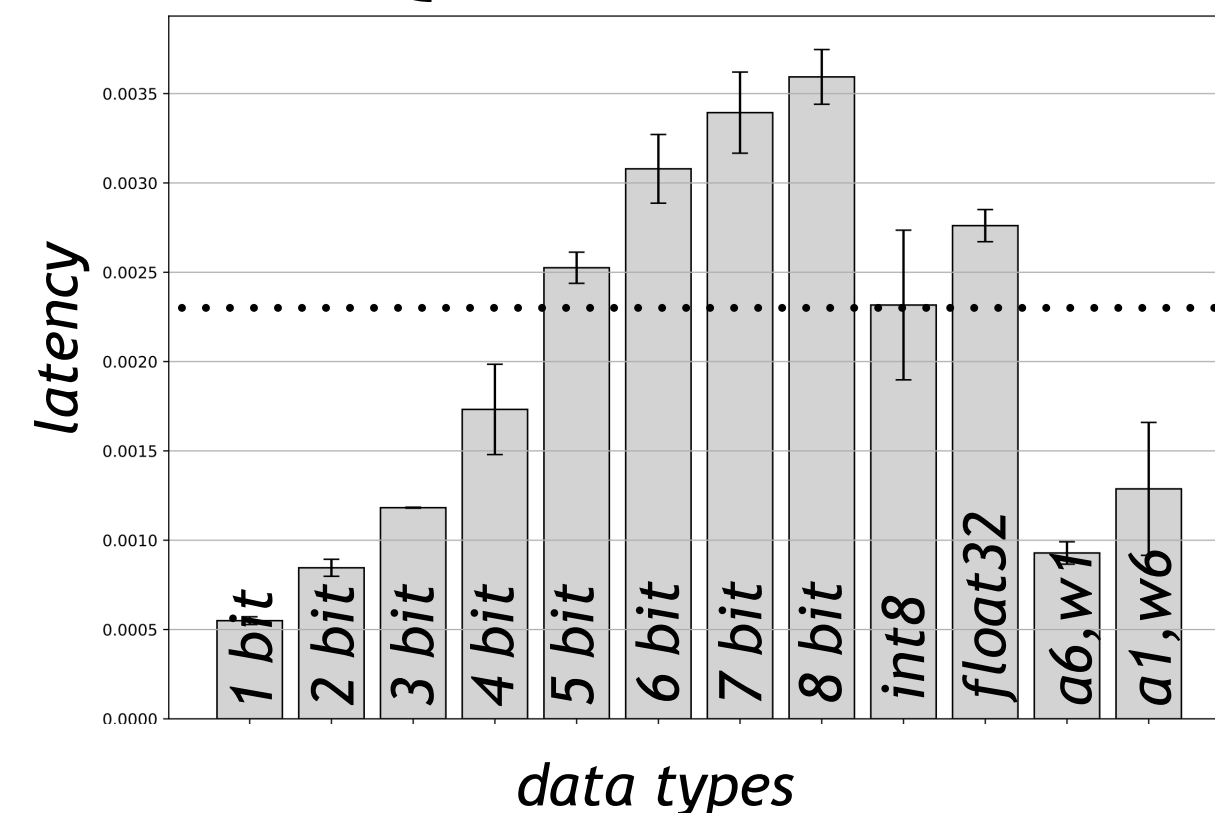
Bender, *et al.* 2020. Can Weight Sharing
Outperform Random Architecture Search?
An Investigation With TuNAS. In *IEEE/CVF
CVPR*.

$$r(P) = acc_{M_P} + \beta \left| \frac{T_{M_P}}{c \cdot T_M} - 1 \right|$$

negative cost exponent

target latency

Quantization



M_{Pe}

Validate M_{Pe}

V_e



Evaluation

Accuracy

(1 ep. retraining)

Latency Benchmark
(quantized operators)



PERFORMANCE OF THE AGENTS

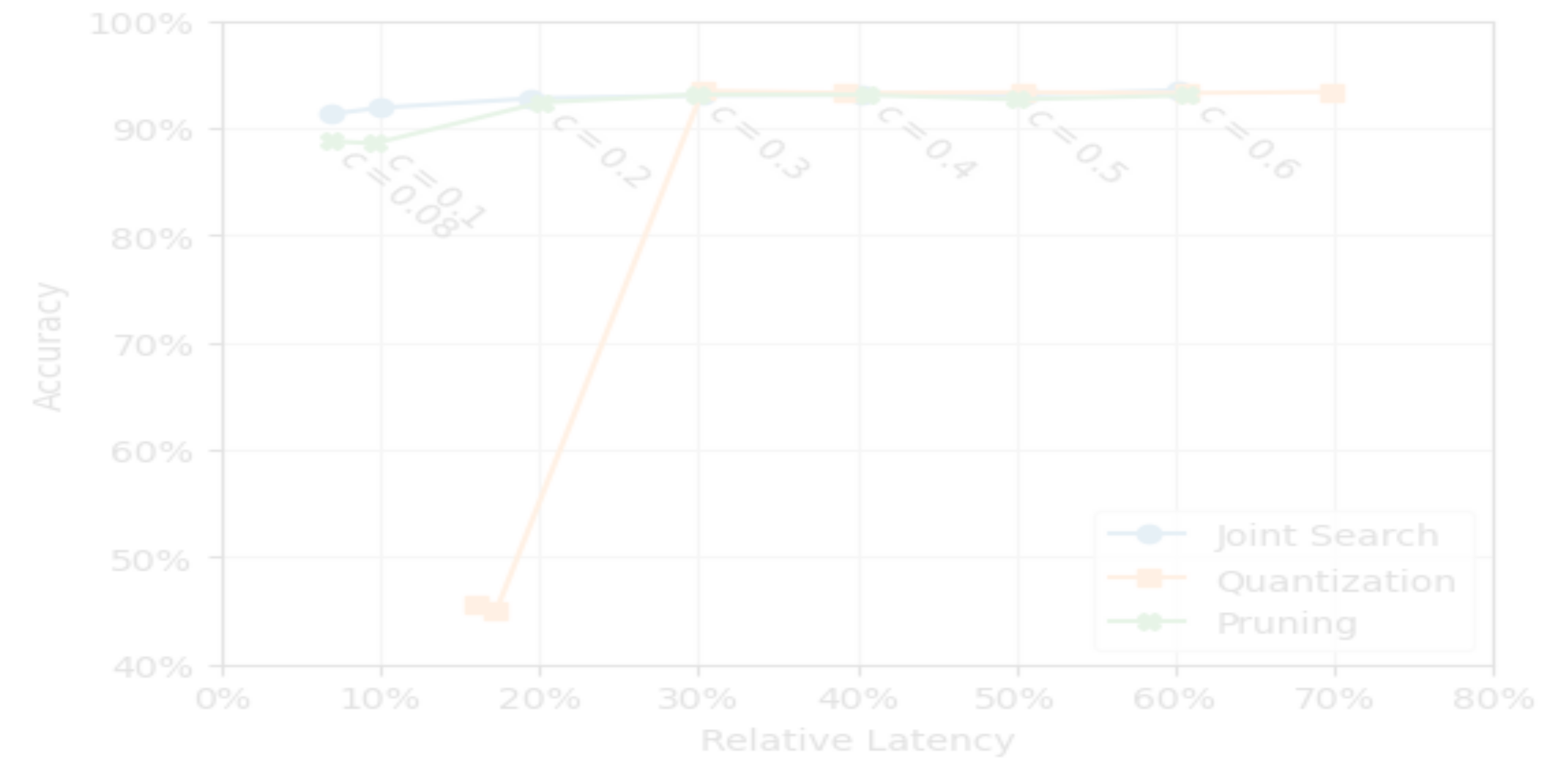
Three different RL agents

Pruning minimizes MACs

Quantization minimizes Bit Operations (BOPs)

Joint approach most balanced

Method	c	MACs	BOPs	Latency	Accuracy
Uncompressed		$4.75 \cdot 10^{10}$	$4.86 \cdot 10^{13}$	330 ms	93.0 %
Pruning Agent		$1.42 \cdot 10^{10}$	$1.45 \cdot 10^{13}$	98 ms	93.0 %
Quantization A.	0.3	$4.75 \cdot 10^{10}$	$8.23 \cdot 10^{11}$	98 ms	92.5 %
Joint Agent		$4.35 \cdot 10^{10}$	$9.42 \cdot 10^{11}$	99 ms	93.2 %
Pruning Agent		$9.24 \cdot 10^9$	$9.45 \cdot 10^{12}$	66 ms	92.4 %
Quantization A.	0.2	$4.75 \cdot 10^{10}$	$4.01 \cdot 10^{11}$	57 ms	45.0 %
Joint Agent		$2.82 \cdot 10^{10}$	$6.74 \cdot 10^{11}$	64 ms	92.8 %



Performance for Various Target Latencies

PERFORMANCE OF THE AGENTS

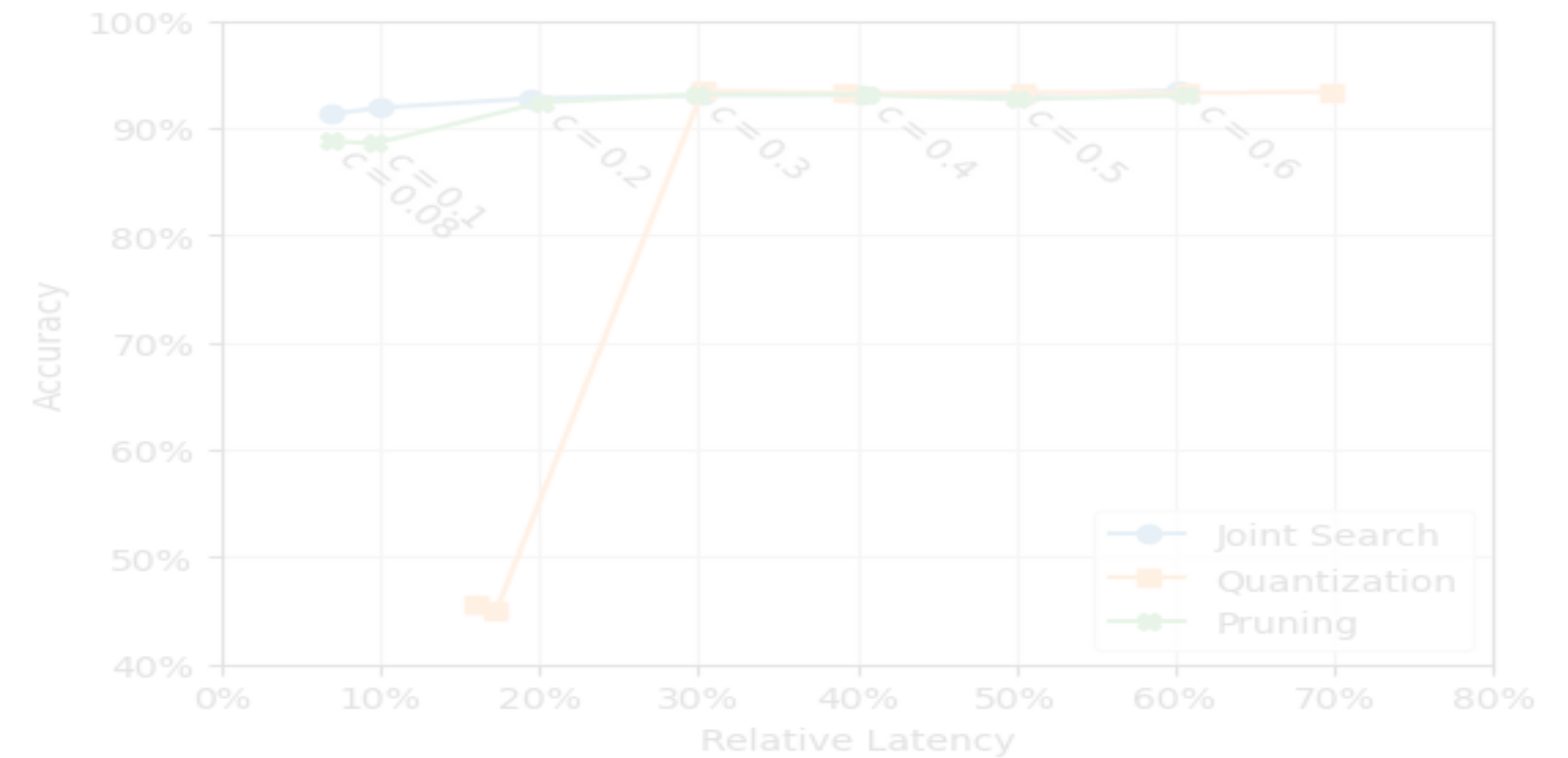
Three different RL agents

Pruning minimizes MACs

Quantization minimizes Bit Operations (BOPs)

Joint approach most balanced

Method	c	MACs	BOPs	Latency	Accuracy
Uncompressed		$4.75 \cdot 10^{10}$	$4.86 \cdot 10^{13}$	330 ms	93.0 %
Pruning Agent		$1.42 \cdot 10^{10}$	$1.45 \cdot 10^{13}$	98 ms	93.0 %
Quantization A.	0.3	$4.75 \cdot 10^{10}$	$8.23 \cdot 10^{11}$	98 ms	92.5 %
Joint Agent		$4.35 \cdot 10^{10}$	$9.42 \cdot 10^{11}$	99 ms	93.2 %
Pruning Agent		$9.24 \cdot 10^9$	$9.45 \cdot 10^{12}$	66 ms	92.4 %
Quantization A.	0.2	$4.75 \cdot 10^{10}$	$4.01 \cdot 10^{11}$	57 ms	45.0 %
Joint Agent		$2.82 \cdot 10^{10}$	$6.74 \cdot 10^{11}$	64 ms	92.8 %



Performance for Various Target Latencies

PERFORMANCE OF THE AGENTS

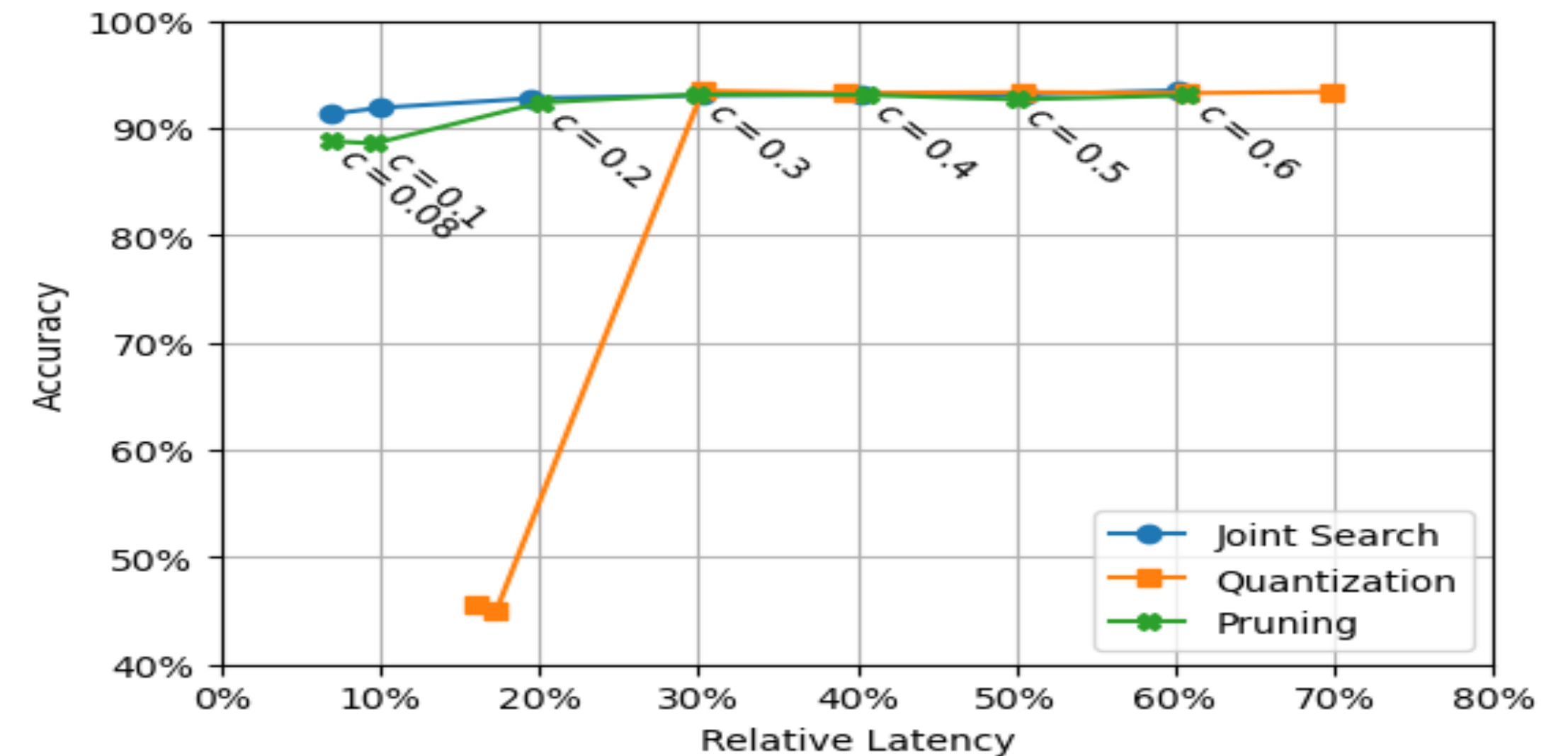
Three different RL agents

Pruning minimizes MACs

Quantization minimizes Bit Operations (BOPs)

Joint approach most balanced

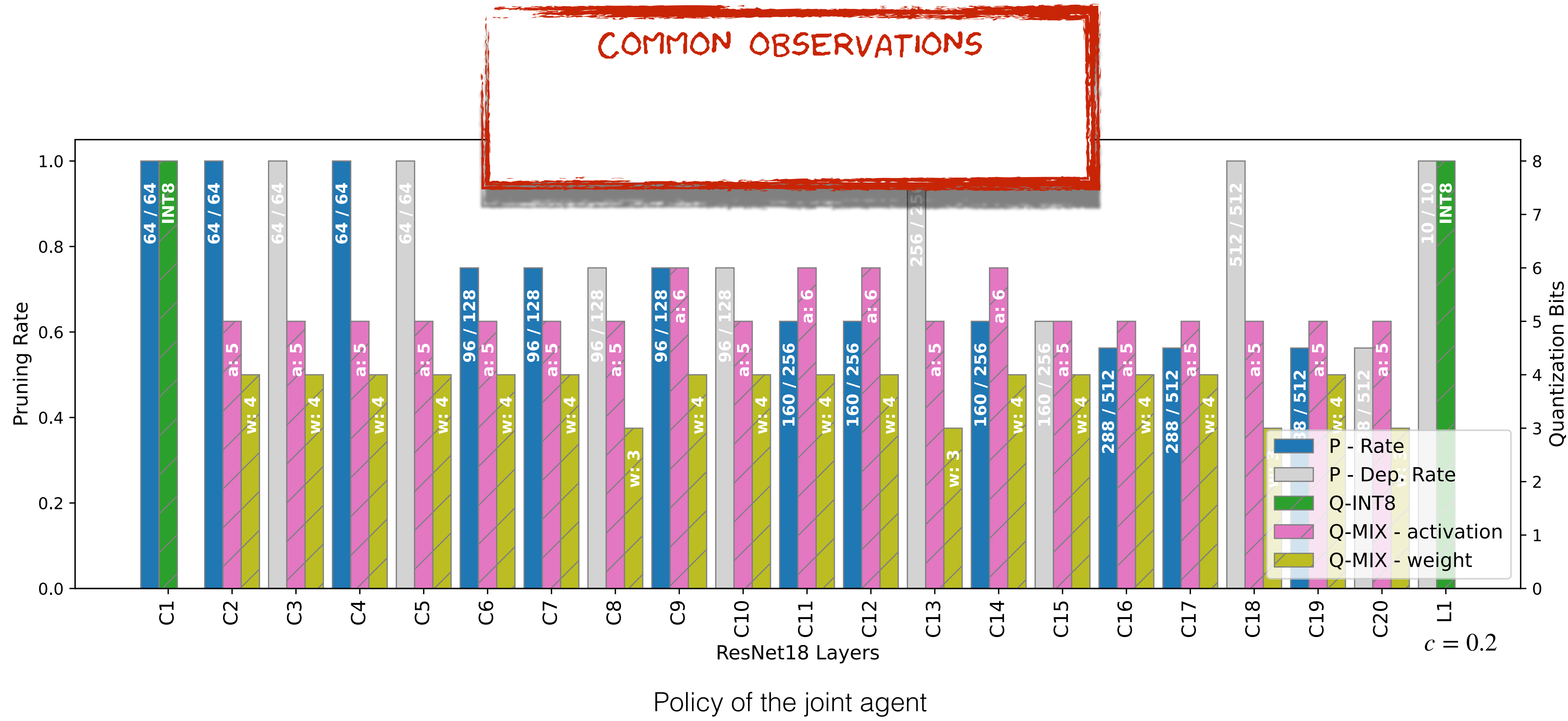
Method	c	MACs	BOPs	Latency	Accuracy
Uncompressed		$4.75 \cdot 10^{10}$	$4.86 \cdot 10^{13}$	330 ms	93.0 %
Pruning Agent		$1.42 \cdot 10^{10}$	$1.45 \cdot 10^{13}$	98 ms	93.0 %
Quantization A.	0.3	$4.75 \cdot 10^{10}$	$8.23 \cdot 10^{11}$	98 ms	92.5 %
Joint Agent		$4.35 \cdot 10^{10}$	$9.42 \cdot 10^{11}$	99 ms	93.2 %
Pruning Agent		$9.24 \cdot 10^9$	$9.45 \cdot 10^{12}$	66 ms	92.4 %
Quantization A.	0.2	$4.75 \cdot 10^{10}$	$4.01 \cdot 10^{11}$	57 ms	45.0 %
Joint Agent		$2.82 \cdot 10^{10}$	$6.74 \cdot 10^{11}$	64 ms	92.8 %



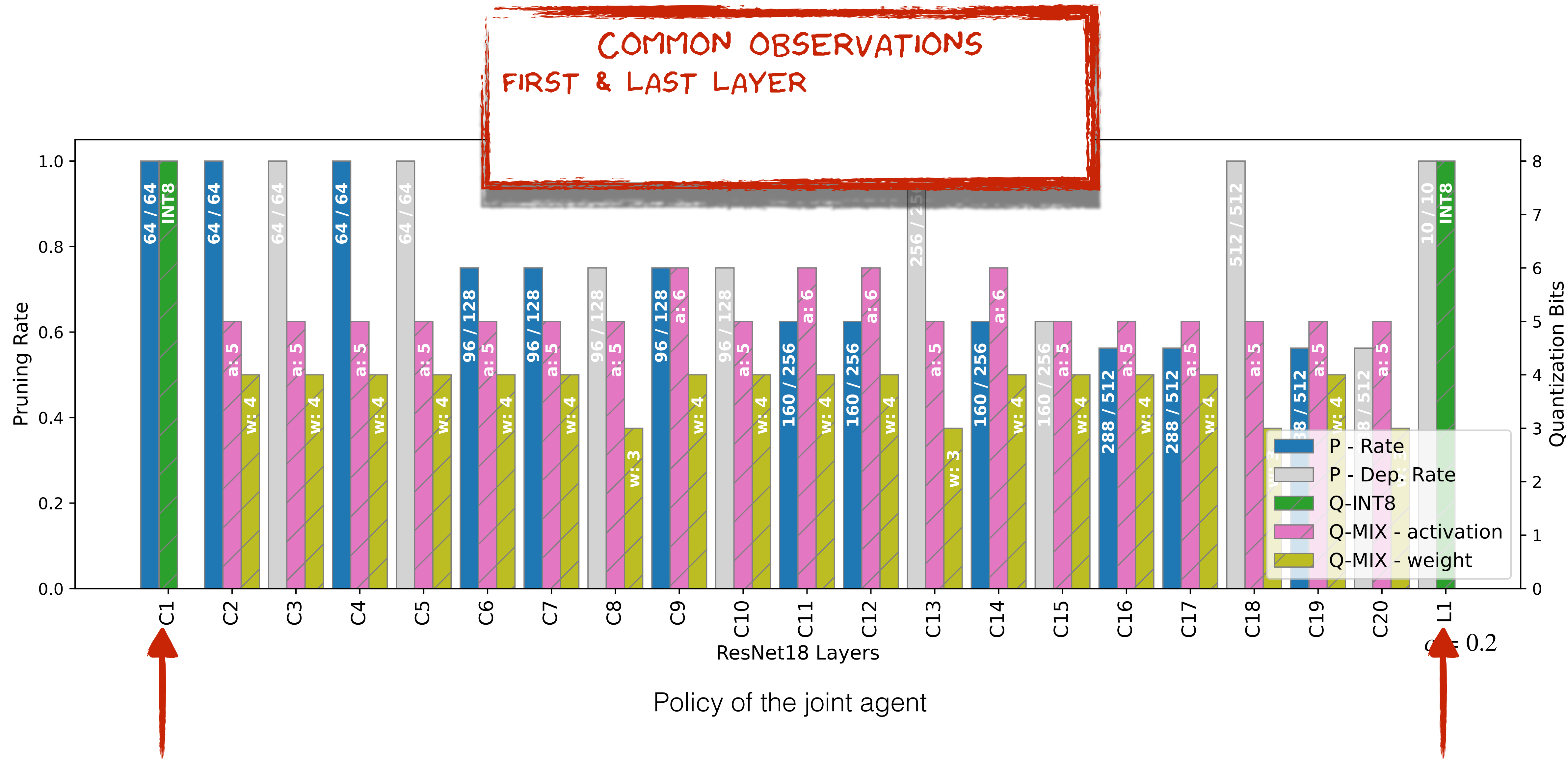
Performance for Various Target Latencies

**CHALLENGING TARGET LATENCIES
FORCE THE AGENT TOWARDS
EXTREME COMPRESSION RESULTING
IN ACCURACY DROPS**

CAN WE UNDERSTAND A PREDICTED POLICY?



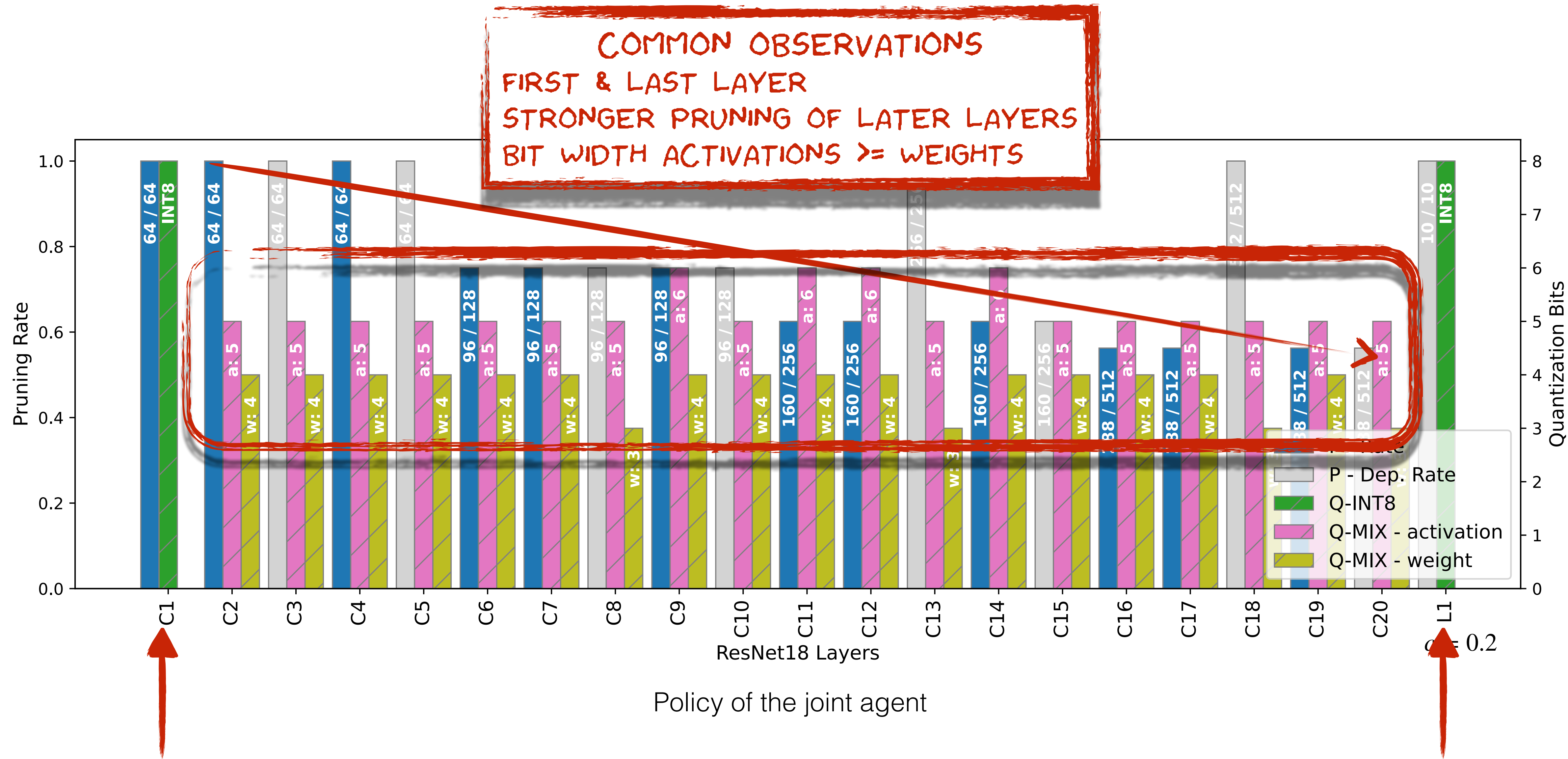
CAN WE UNDERSTAND A PREDICTED POLICY?



CAN WE UNDERSTAND A PREDICTED POLICY?



CAN WE UNDERSTAND A PREDICTED POLICY?

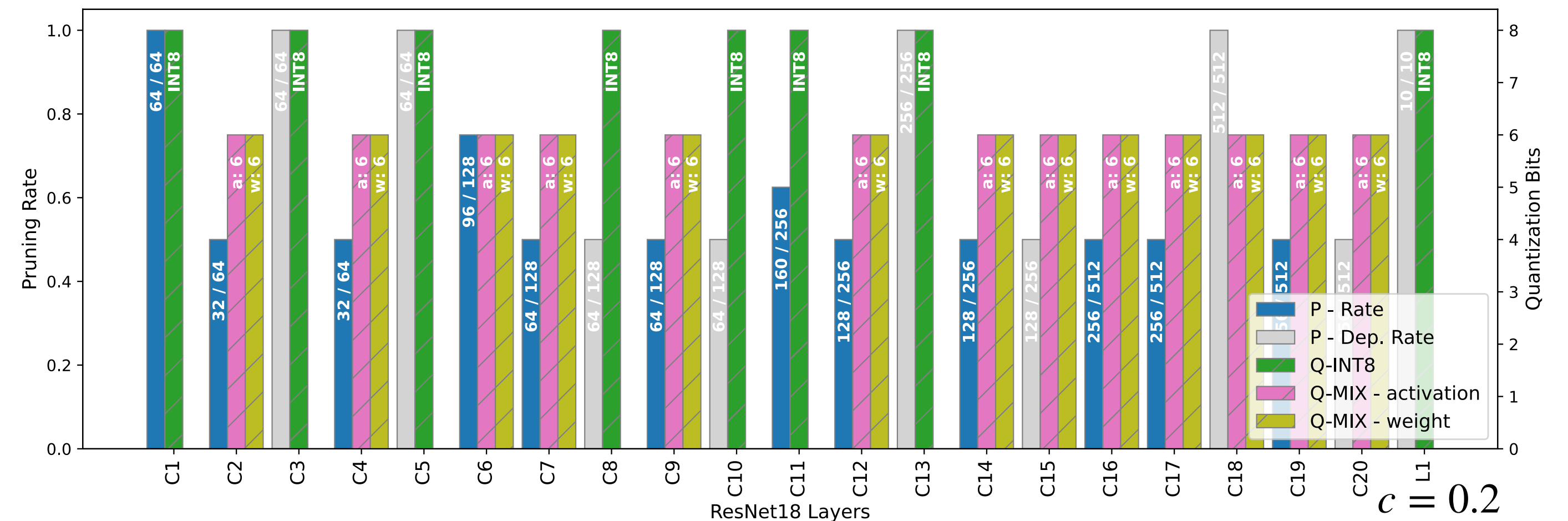


IS THE SENSITIVITY ANALYSIS NECESSARY?

Without sensitivity analysis

Homogeneous pruning

Minimal quantization

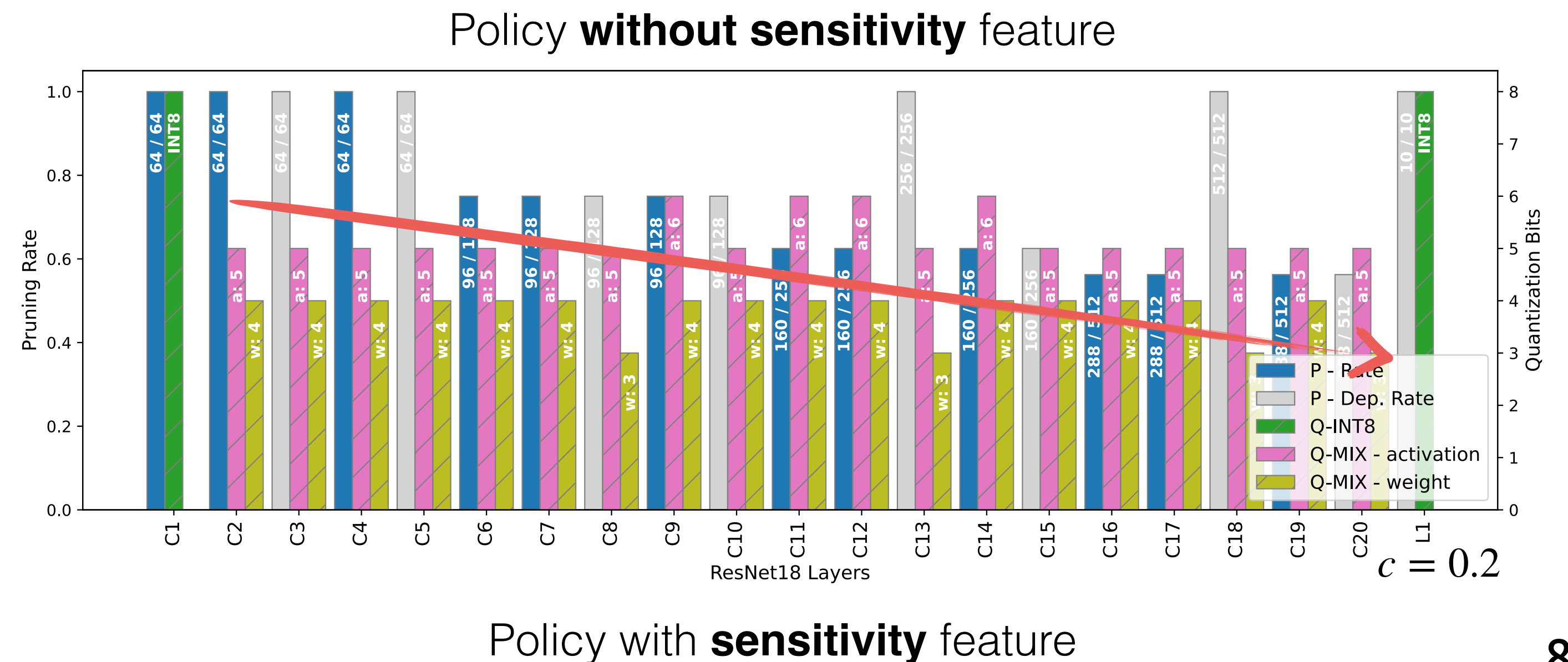


With sensitivity information

Heterogeneous compression

Activations $\neq$ weight precision

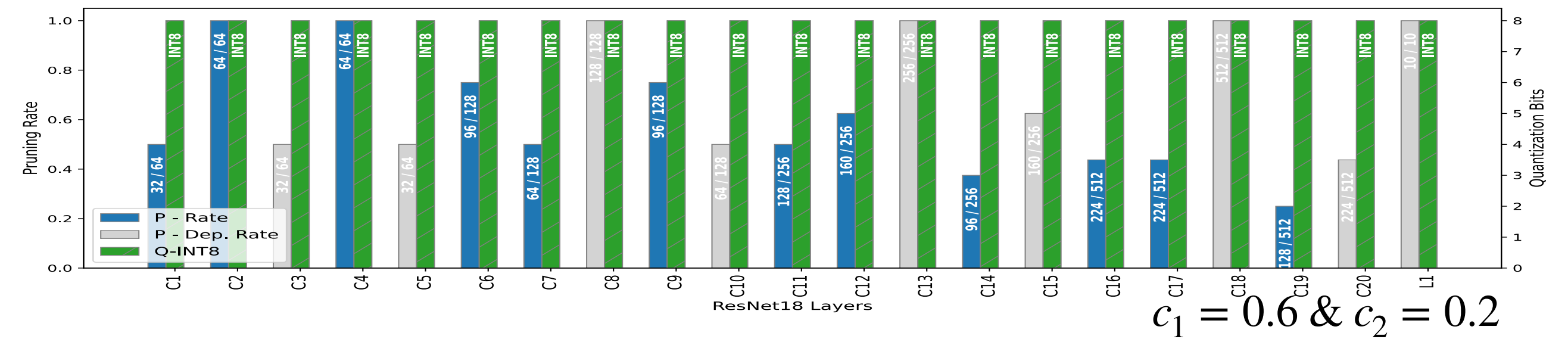
Stronger compression in the last layers



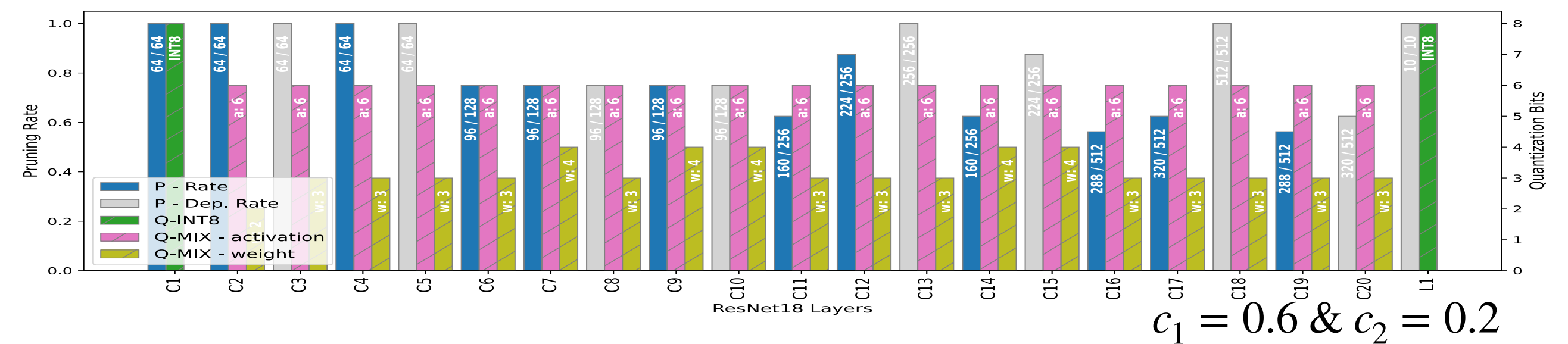
IS A CONCURRENT JOINT SEARCH BENEFICIAL?

Subsequent searches use
latter compression more
aggressively

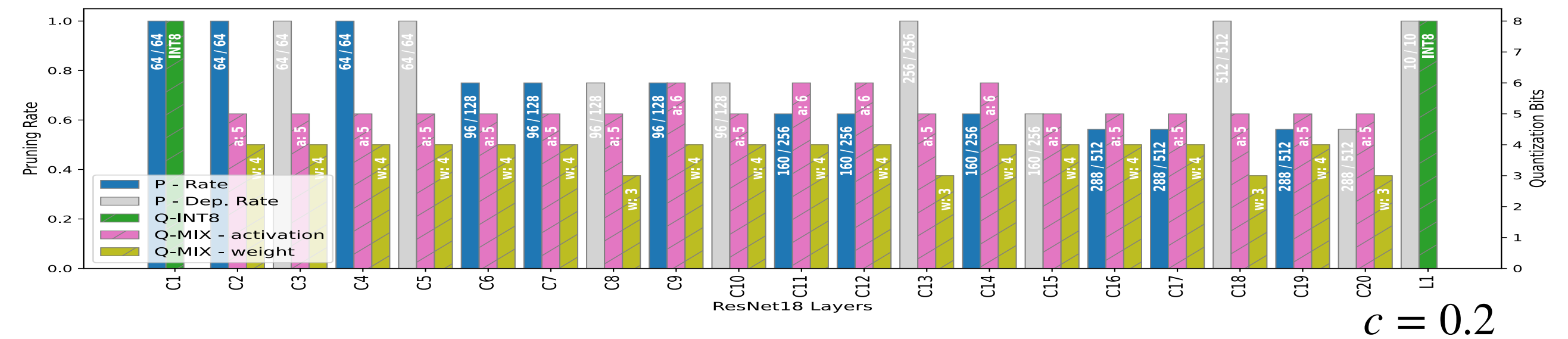
Joint search is more
balanced



First quantization, then pruning search



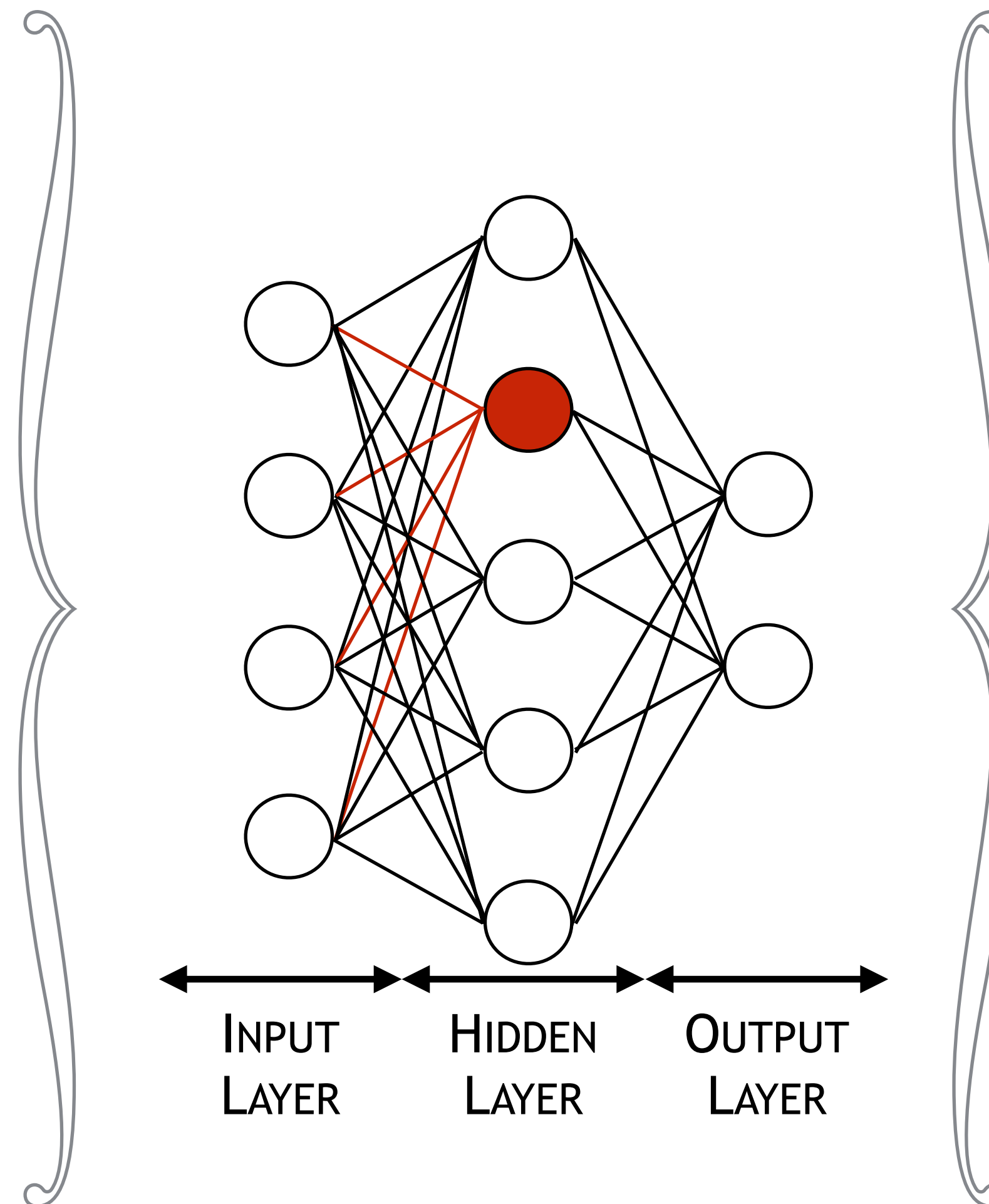
First pruning, then quantization search



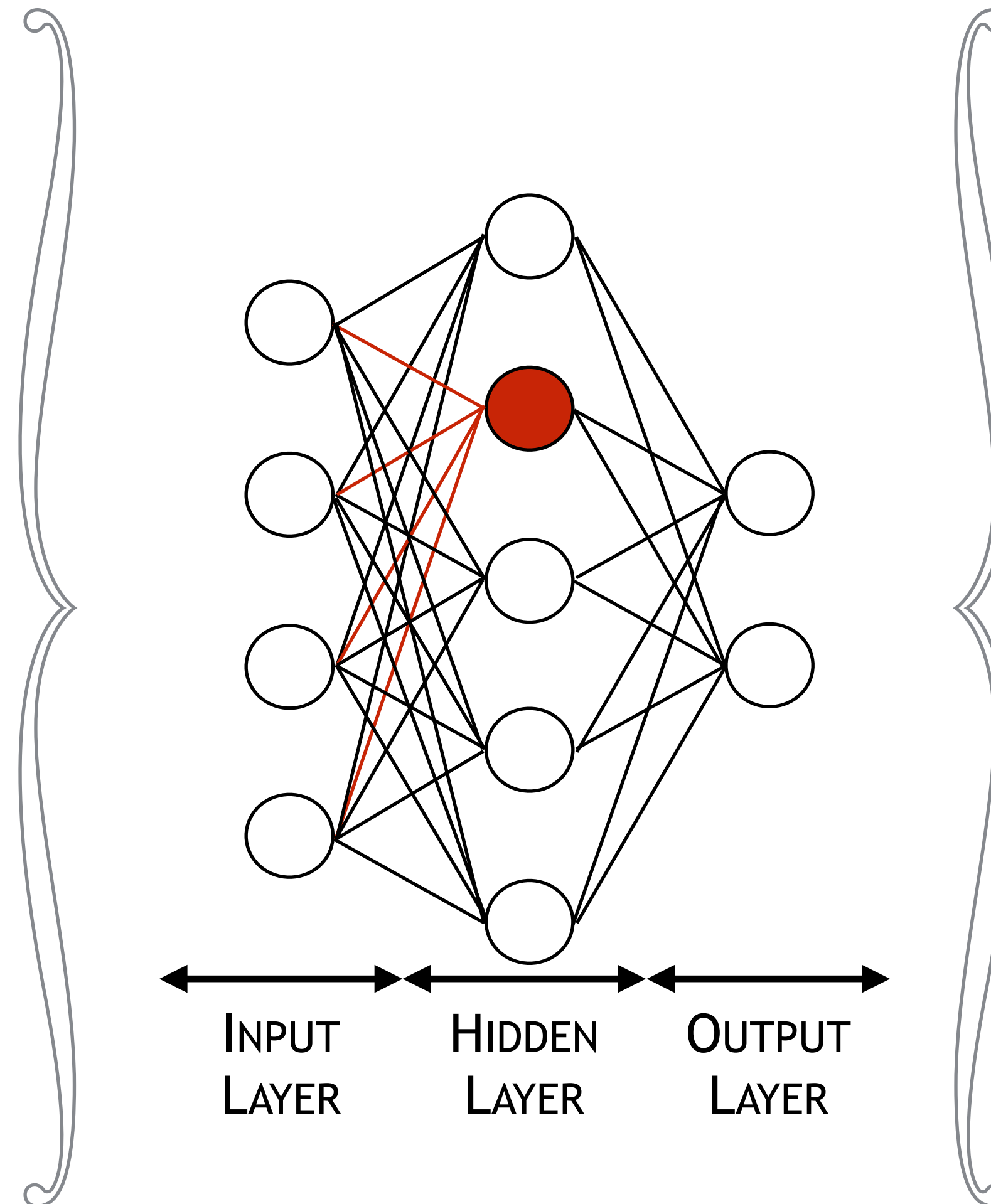
Joint search

BAYESIAN NEURAL NETWORKS

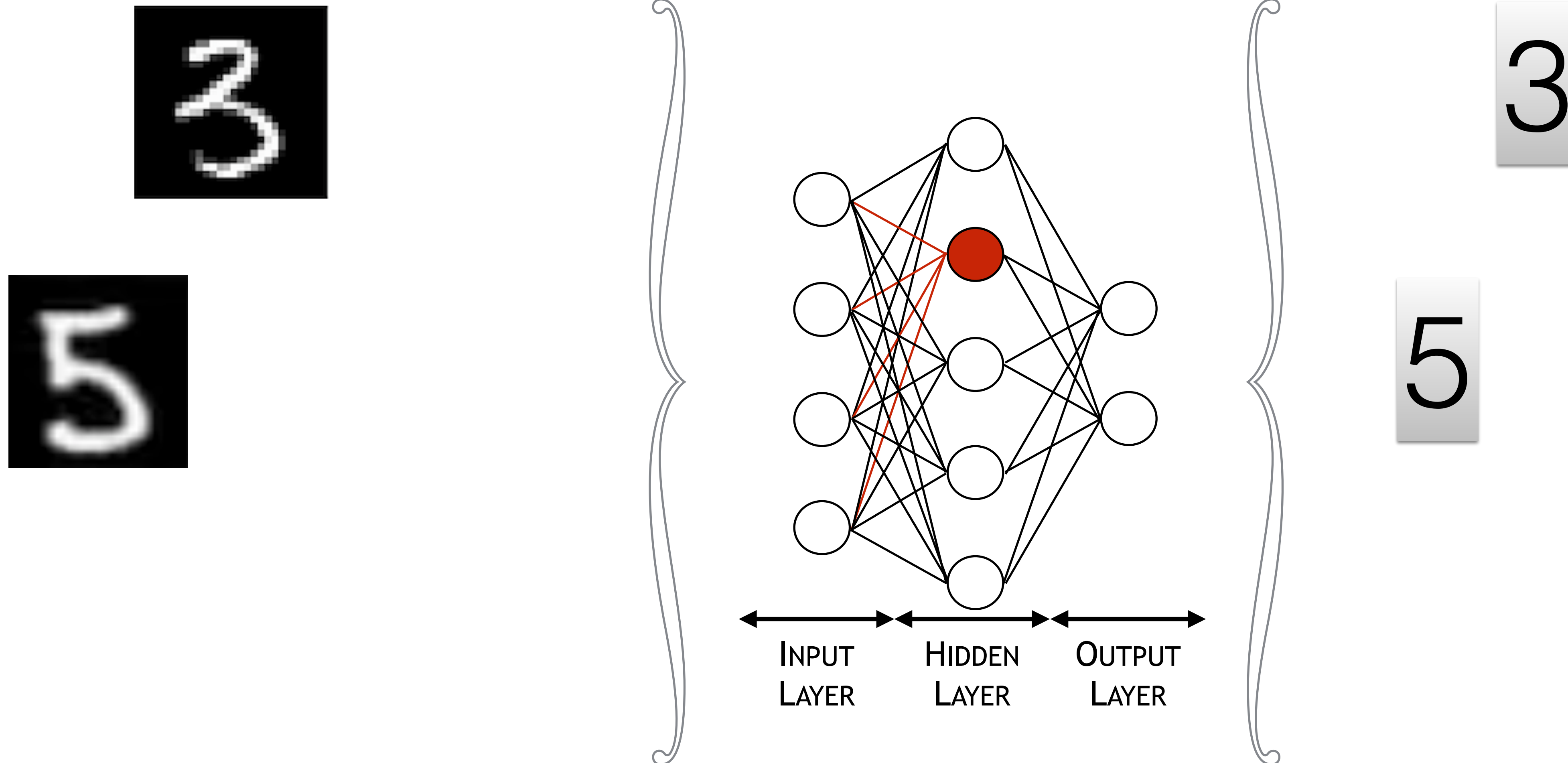
OUT-OF-DOMAIN DATA



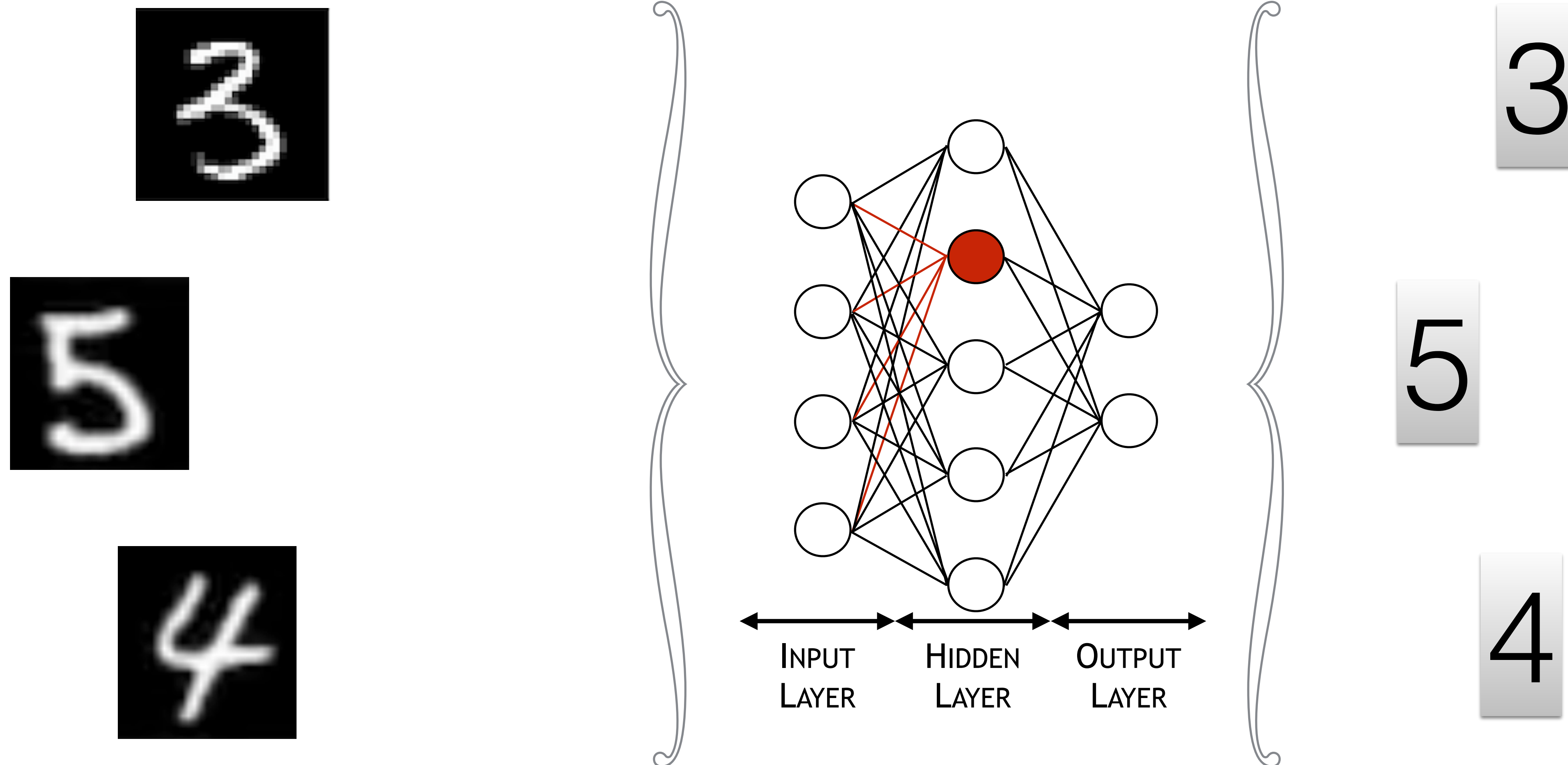
OUT-OF-DOMAIN DATA



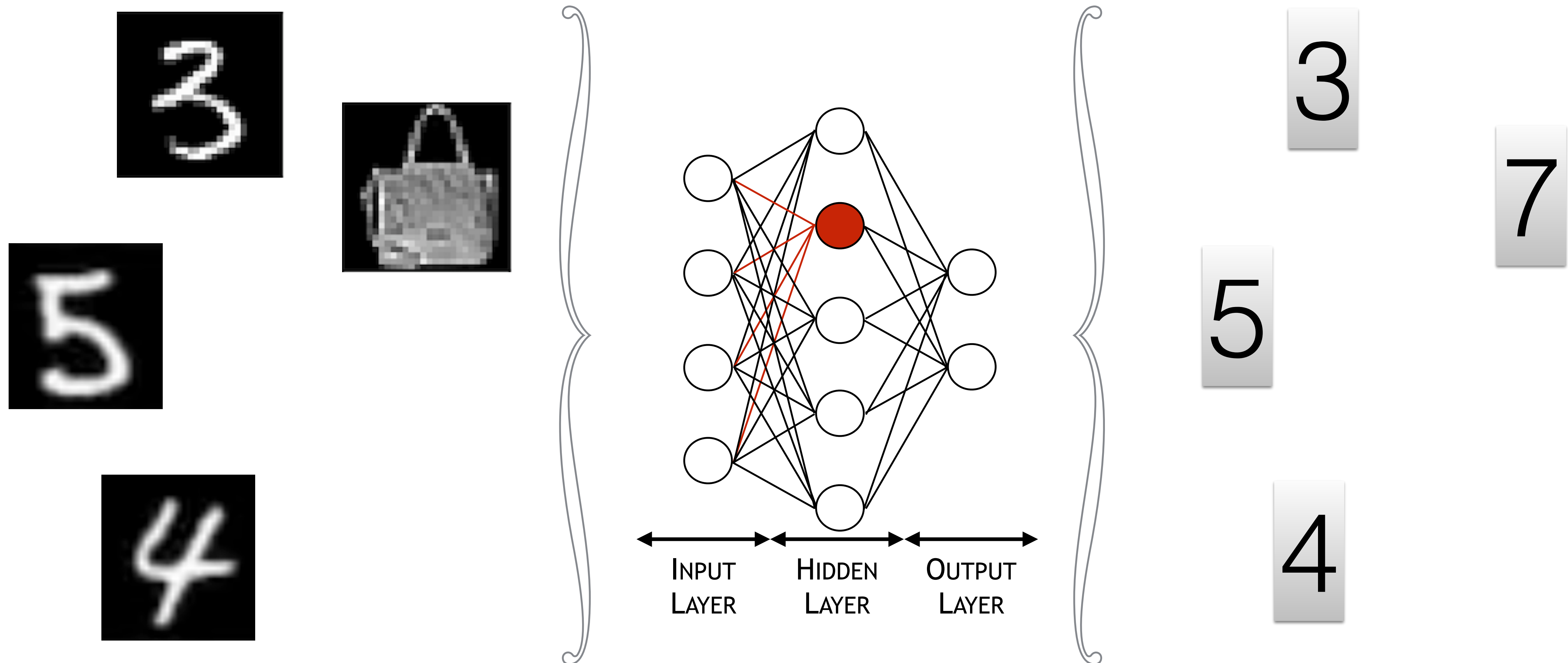
OUT-OF-DOMAIN DATA



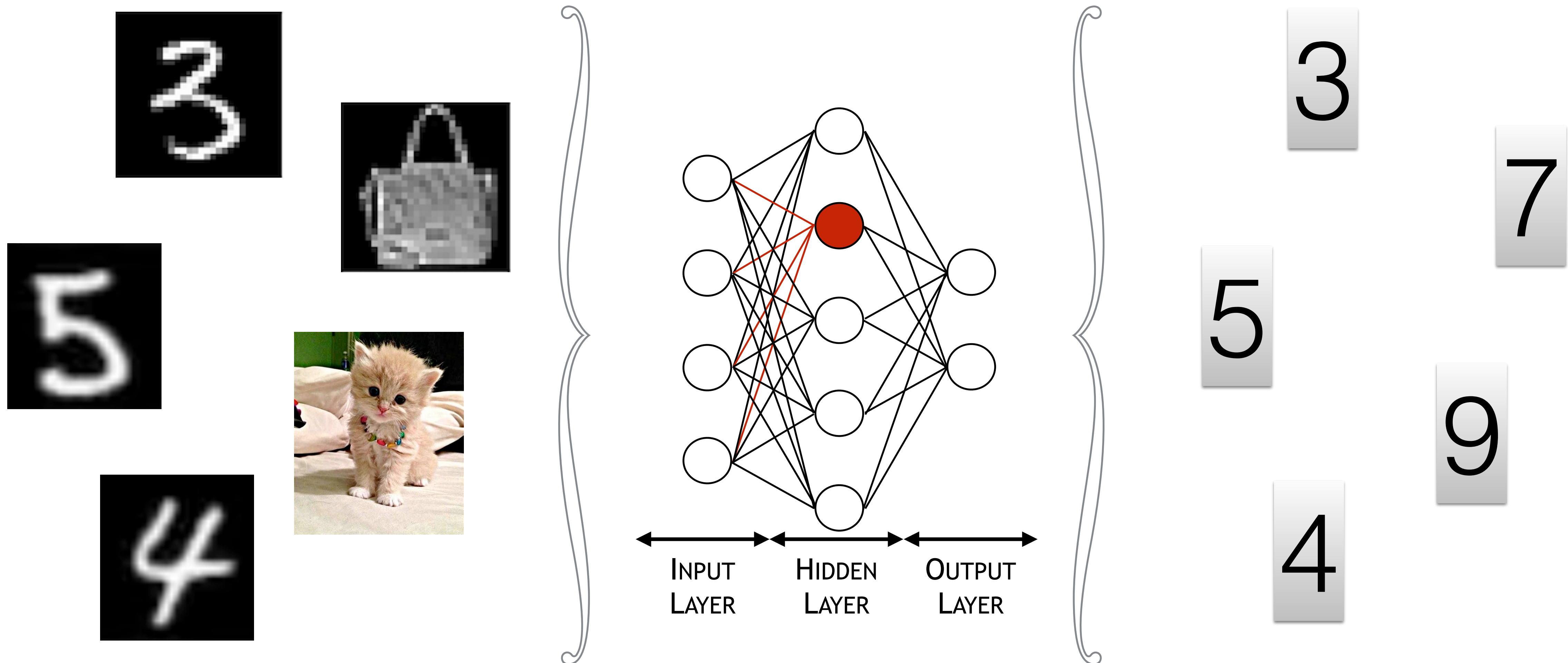
OUT-OF-DOMAIN DATA



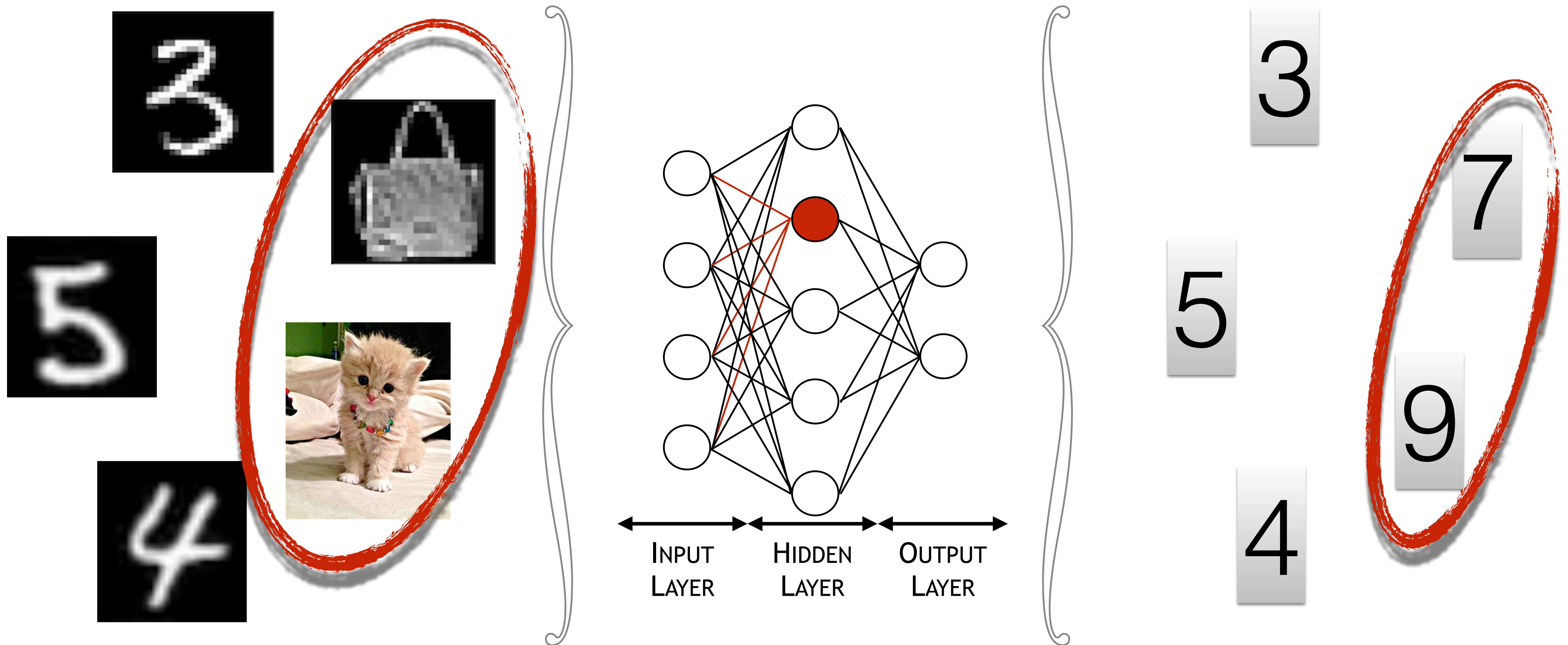
OUT-OF-DOMAIN DATA



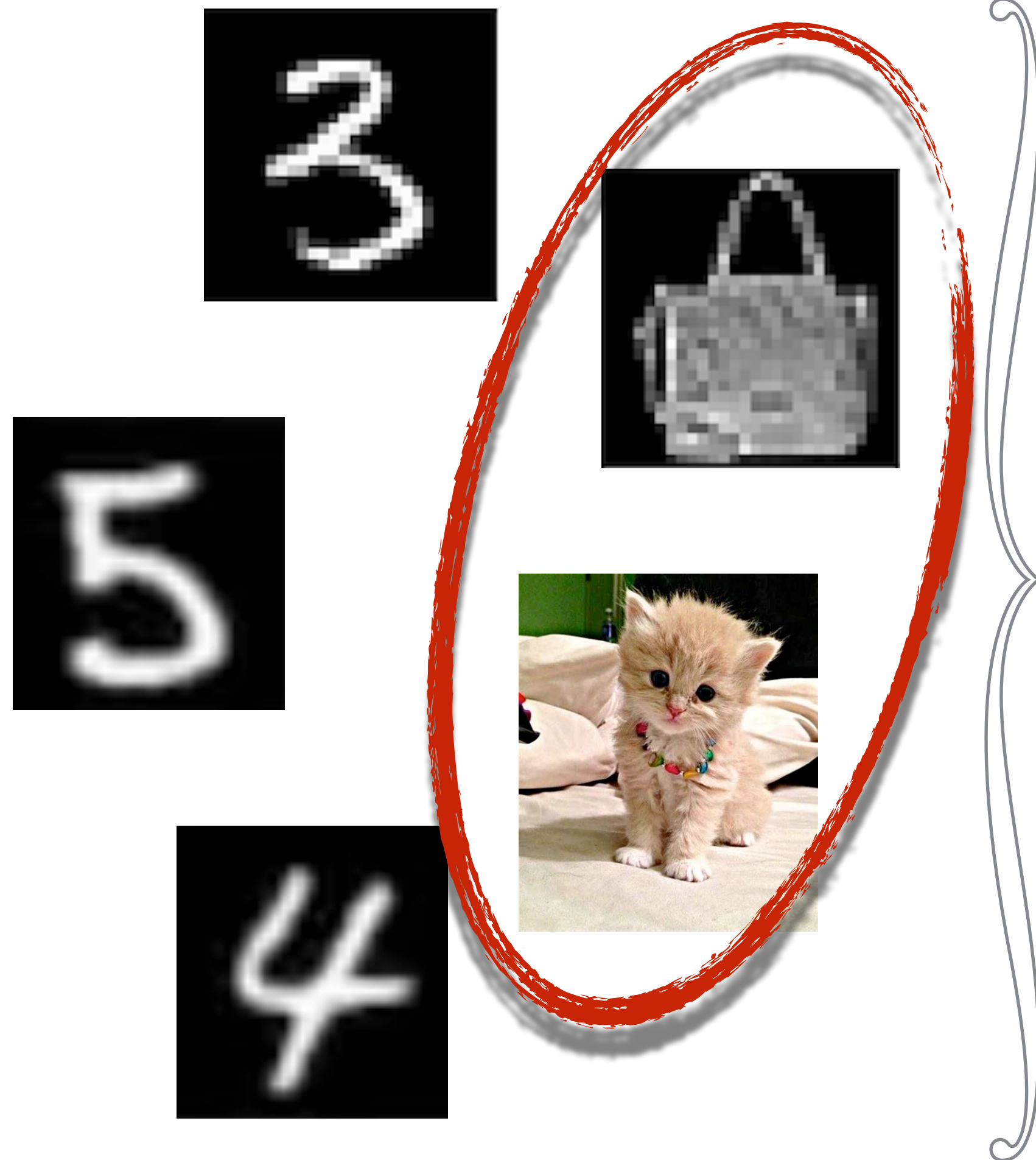
OUT-OF-DOMAIN DATA



OUT-OF-DOMAIN DATA



OUT-OF-DOMAIN DATA



Driverless Car Gets Stuck in Wet Concrete in San Francisco

Though driverless cars have not been blamed for any serious injuries or crashes in the city, they have been involved in several jarring episodes.

 Share full article   376



3

7

9

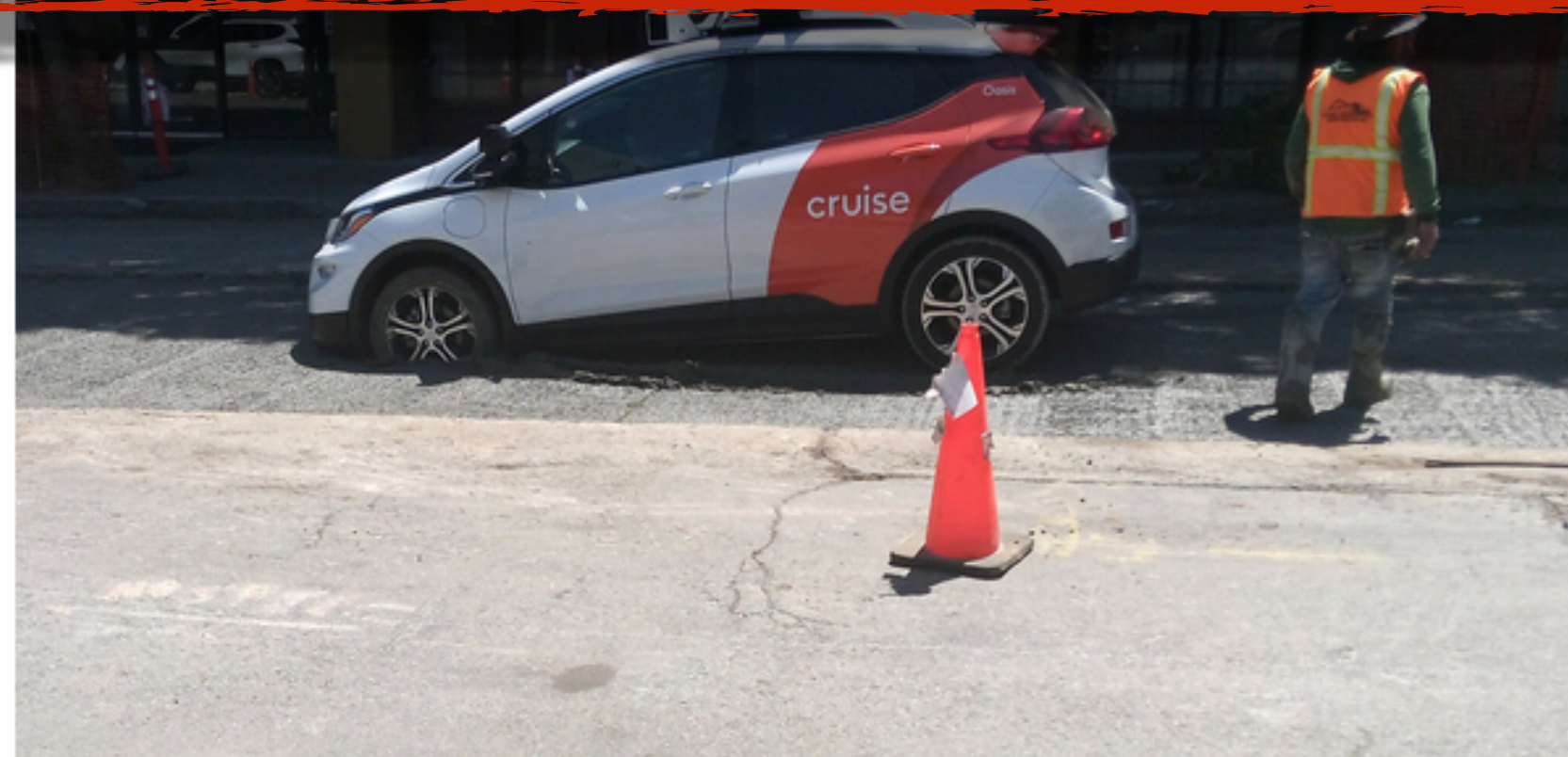
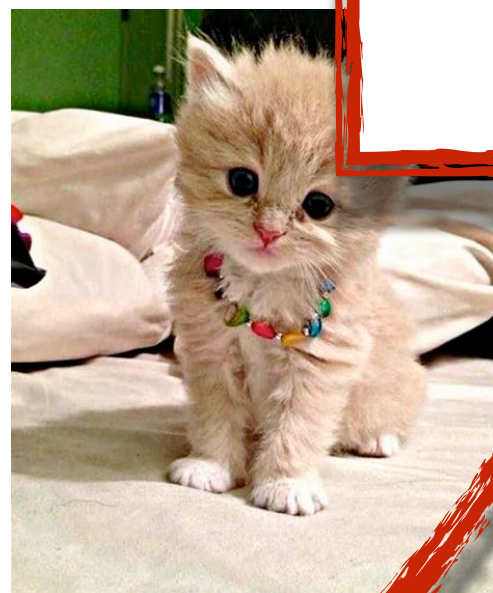
4

OUT-OF-DOMAIN DATA

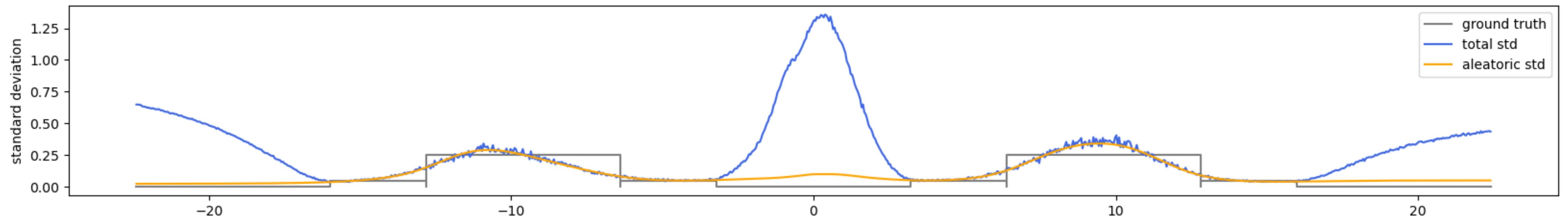
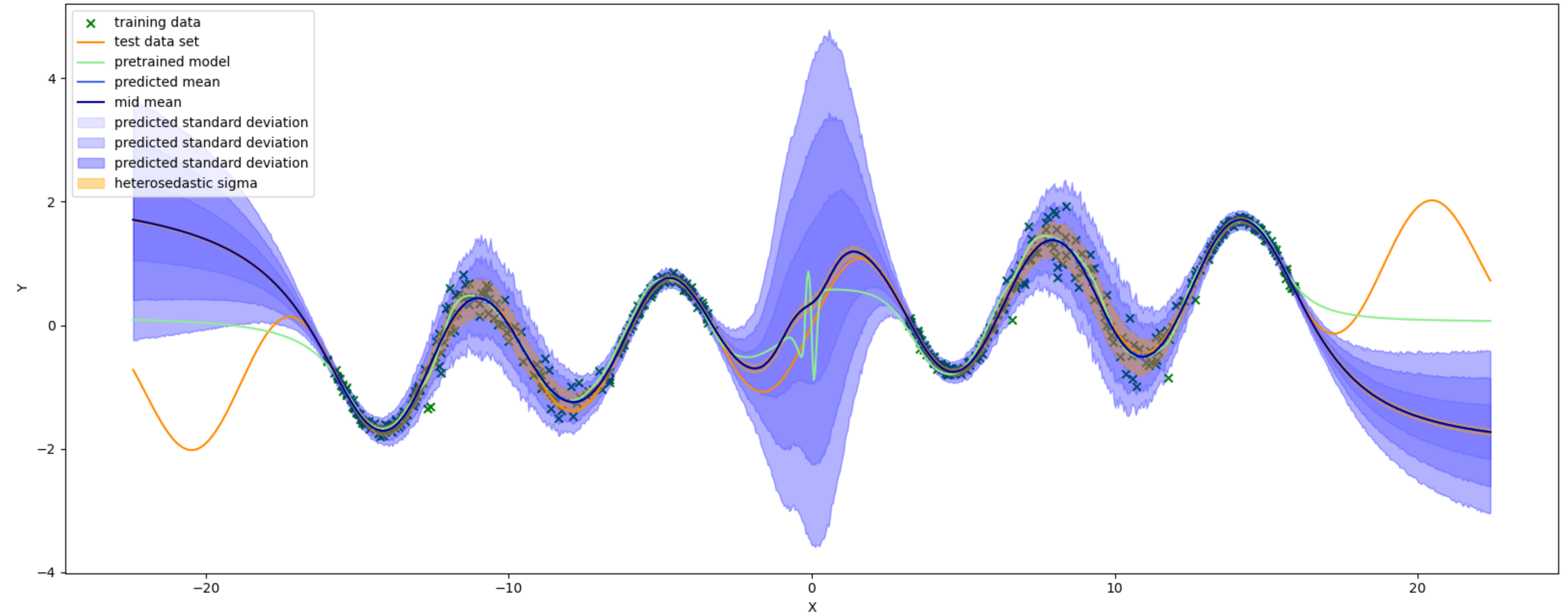
CLASSICAL NEURAL NETWORKS
DON'T HAVE THE ABILITY TO SAY
- I DON'T KNOW -

Driverless Car Gets Stuck in Wet Concrete in San Francisco

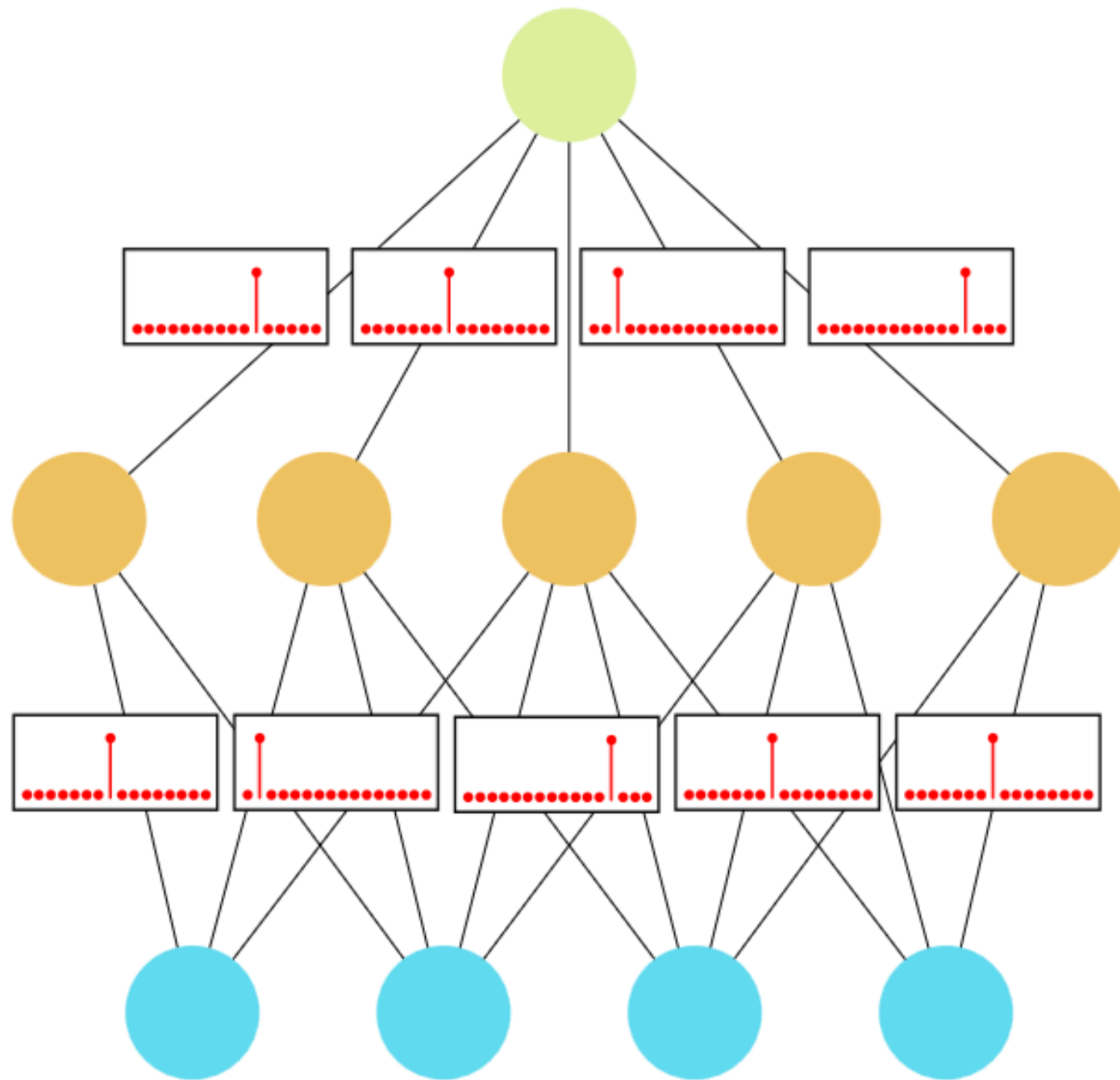
Though driverless cars have not been blamed for any serious injuries or crashes in the city, they have been involved in several



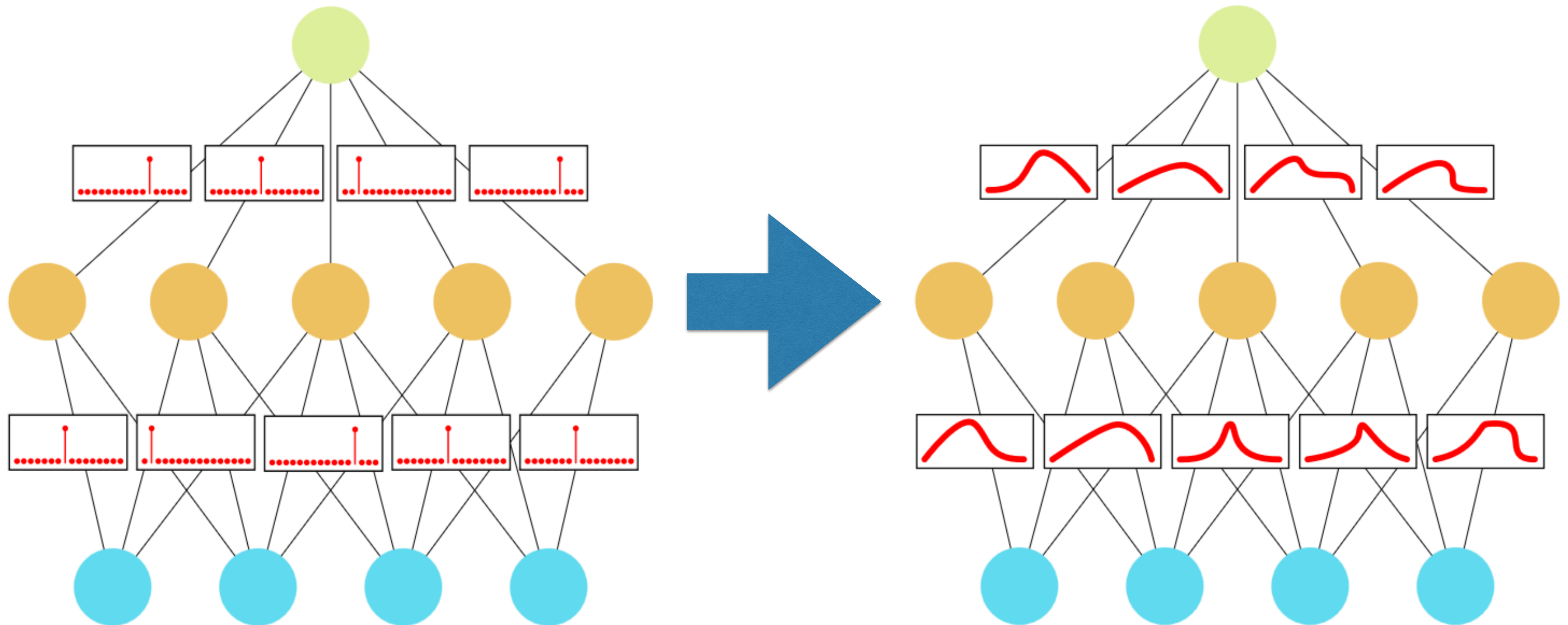
UNCERTAINTY PREDICTION



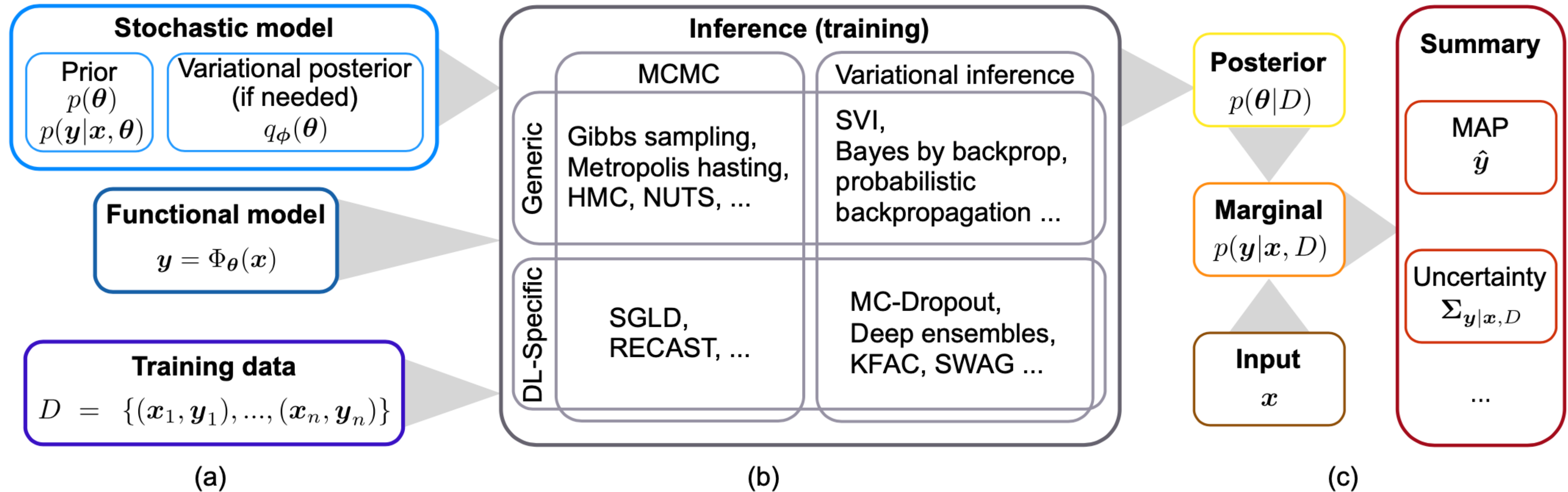
BAYESIAN NEURAL NETWORKS



BAYESIAN NEURAL NETWORKS

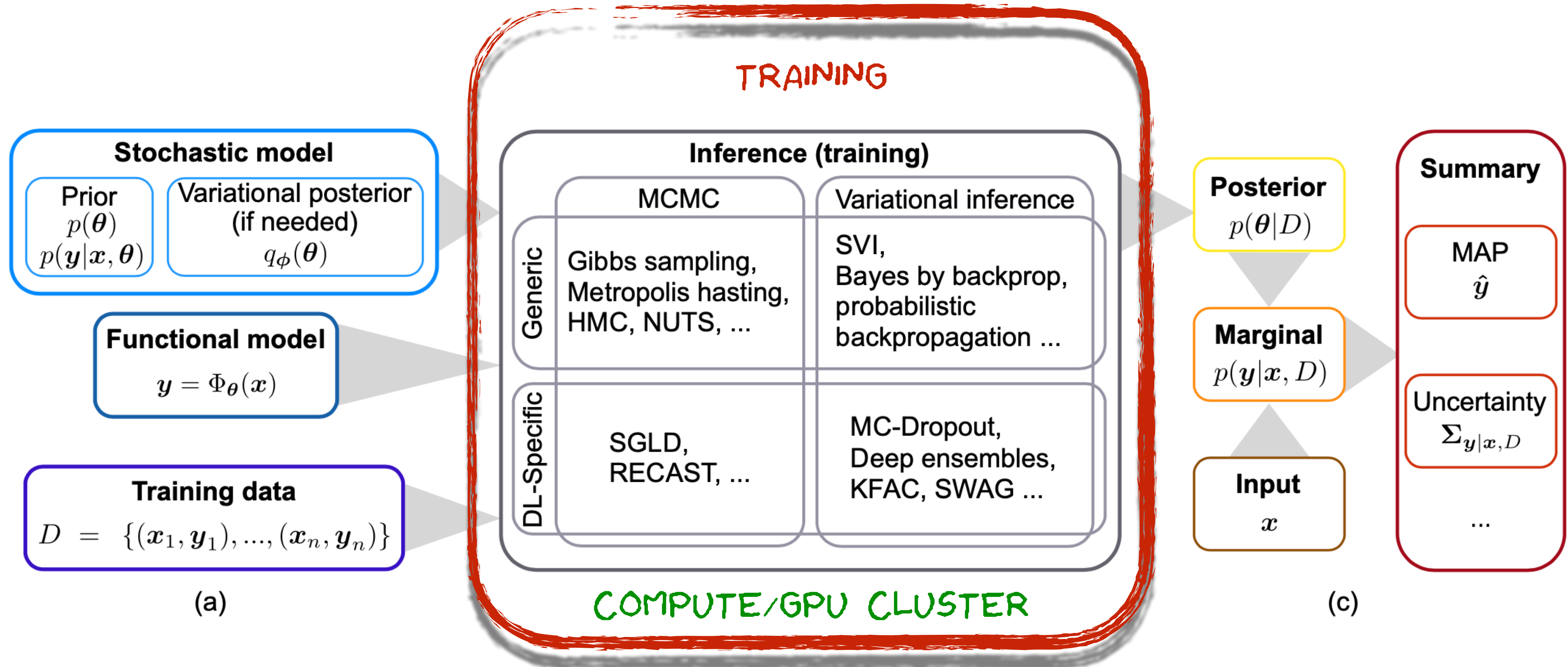


BNN WORKFLOW



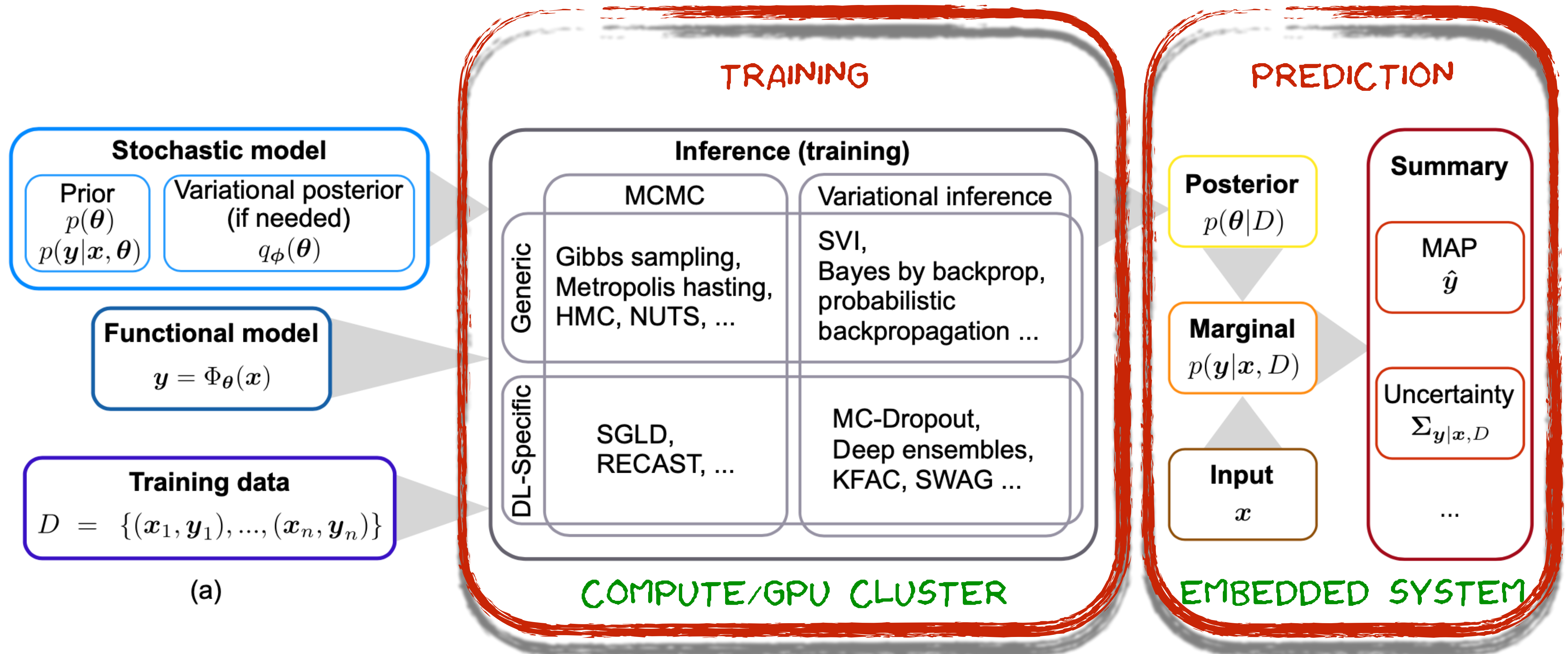
Workflow to design (a), train (b) and use a BNN for predictions (c)

BNN WORKFLOW



Workflow to design (a), train (b) and use a BNN for predictions (c)

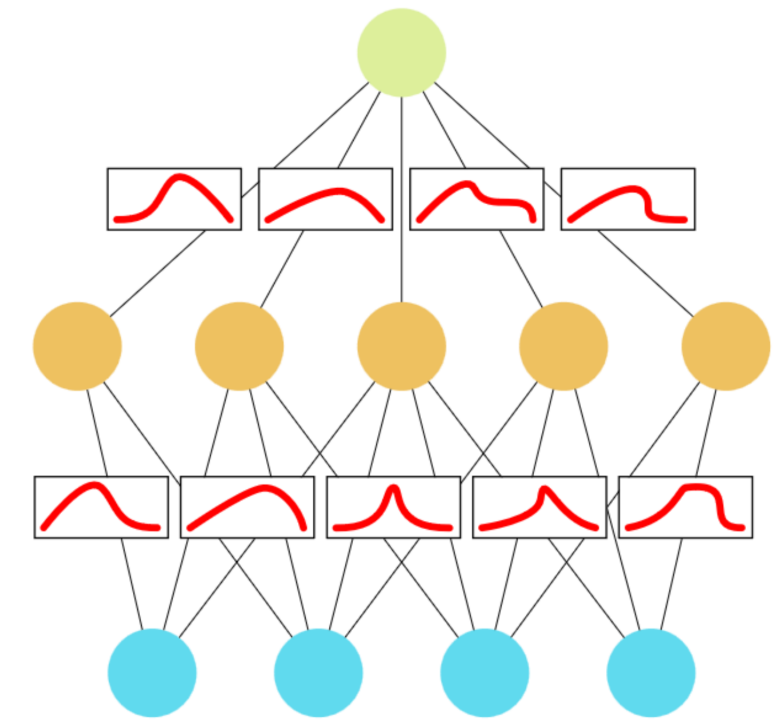
BNN WORKFLOW



Workflow to design (a), train (b) and use a BNN for predictions (c)

BAYESIAN INFERENCE FOR BNNS

Finding the weight distributions



Markov Chain Monte Carlo (MCMC)

Approximate complex distributions of arbitrary shape by drawing many samples

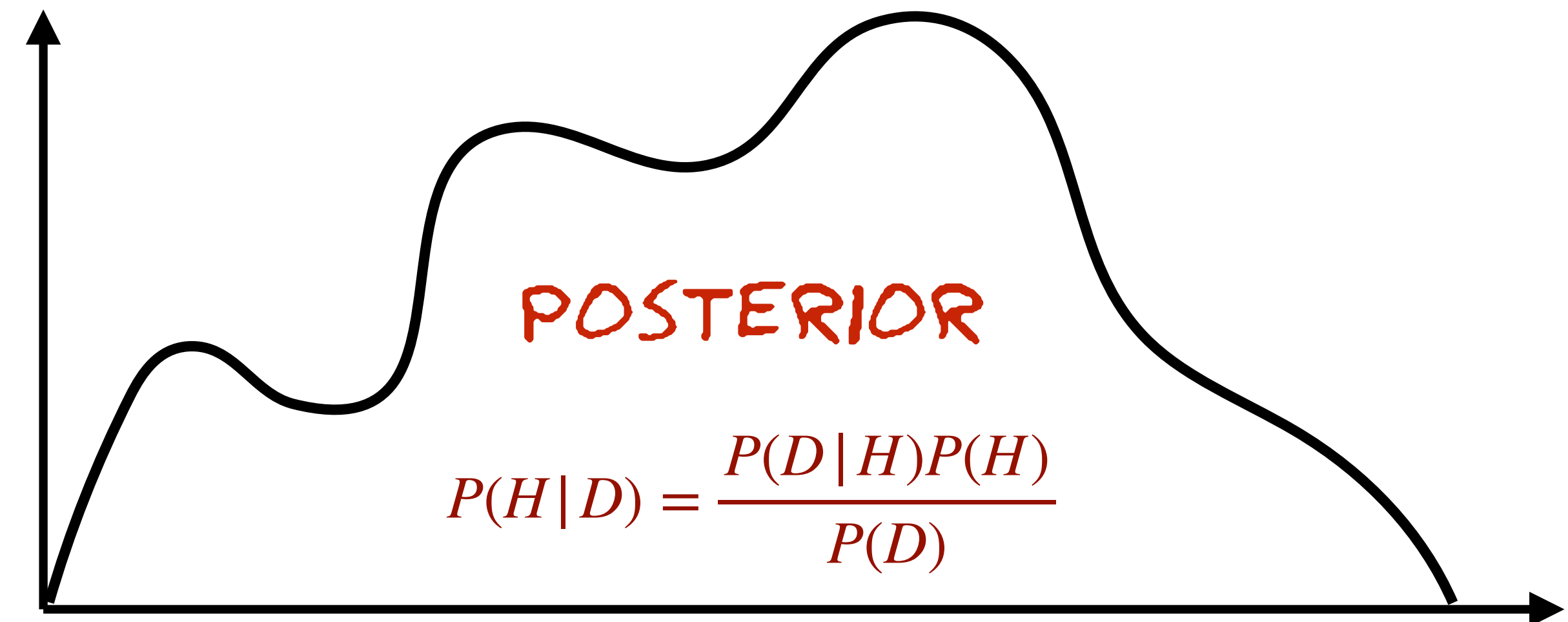
Variational Inference (VI)

Approximate all distributions with simple parametrized (Gaussian) distributions

NN-based approximations

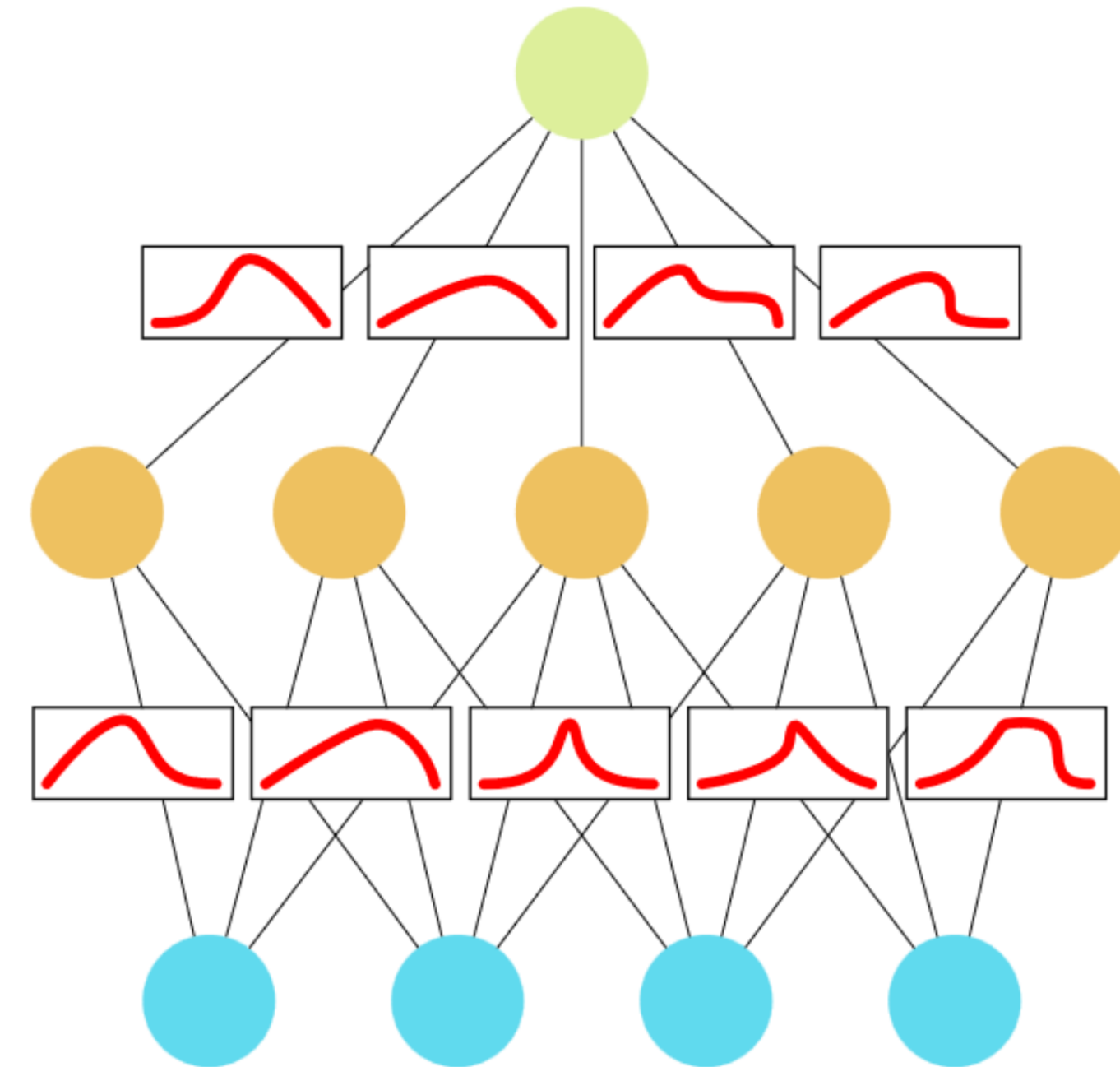
DeepEnsembles (DE)

Monte Carlo Dropout (MCDO)



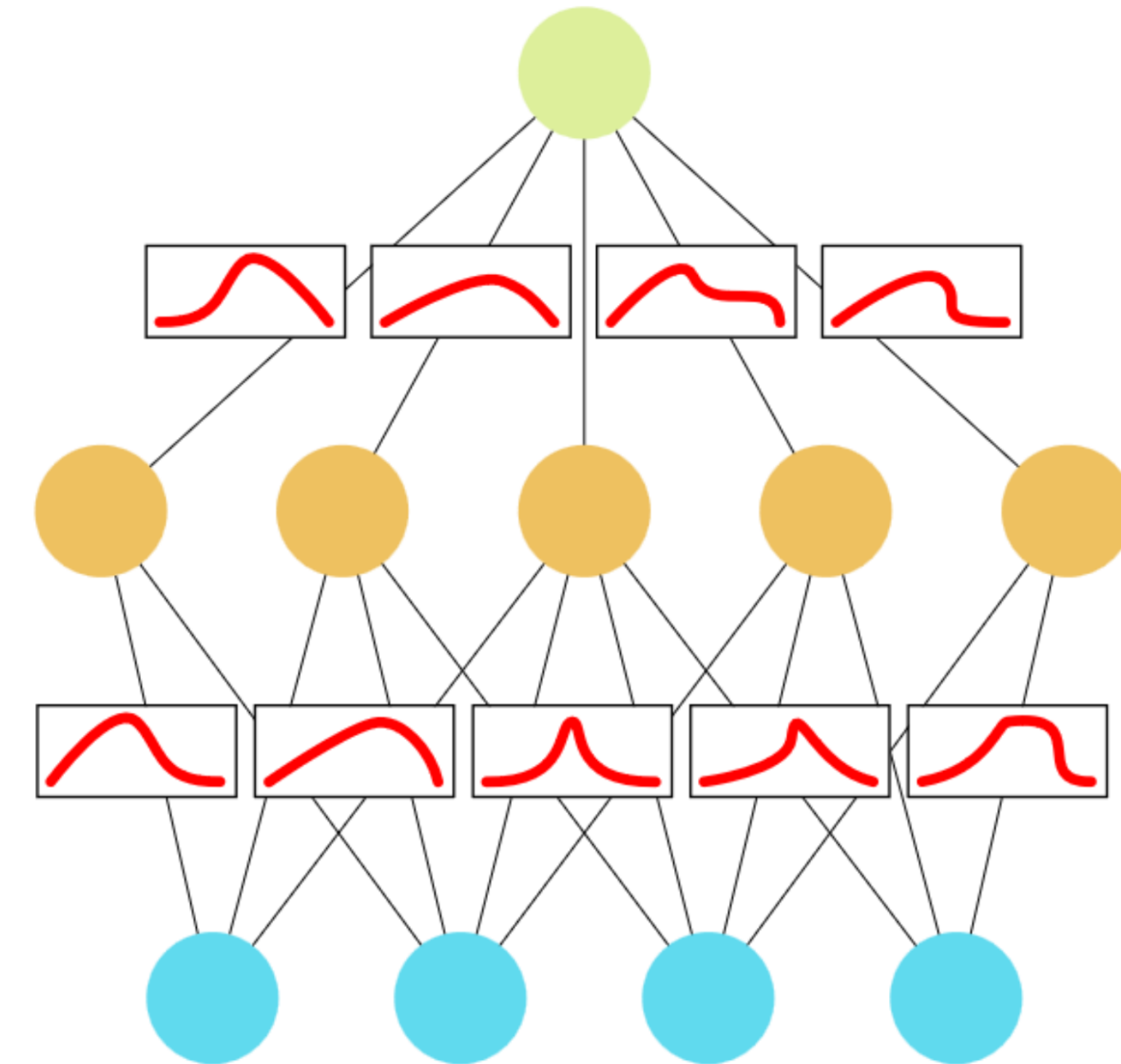
BAYESIAN NEURAL NETWORKS

1. Predict weight distributions with Bayesian Inference



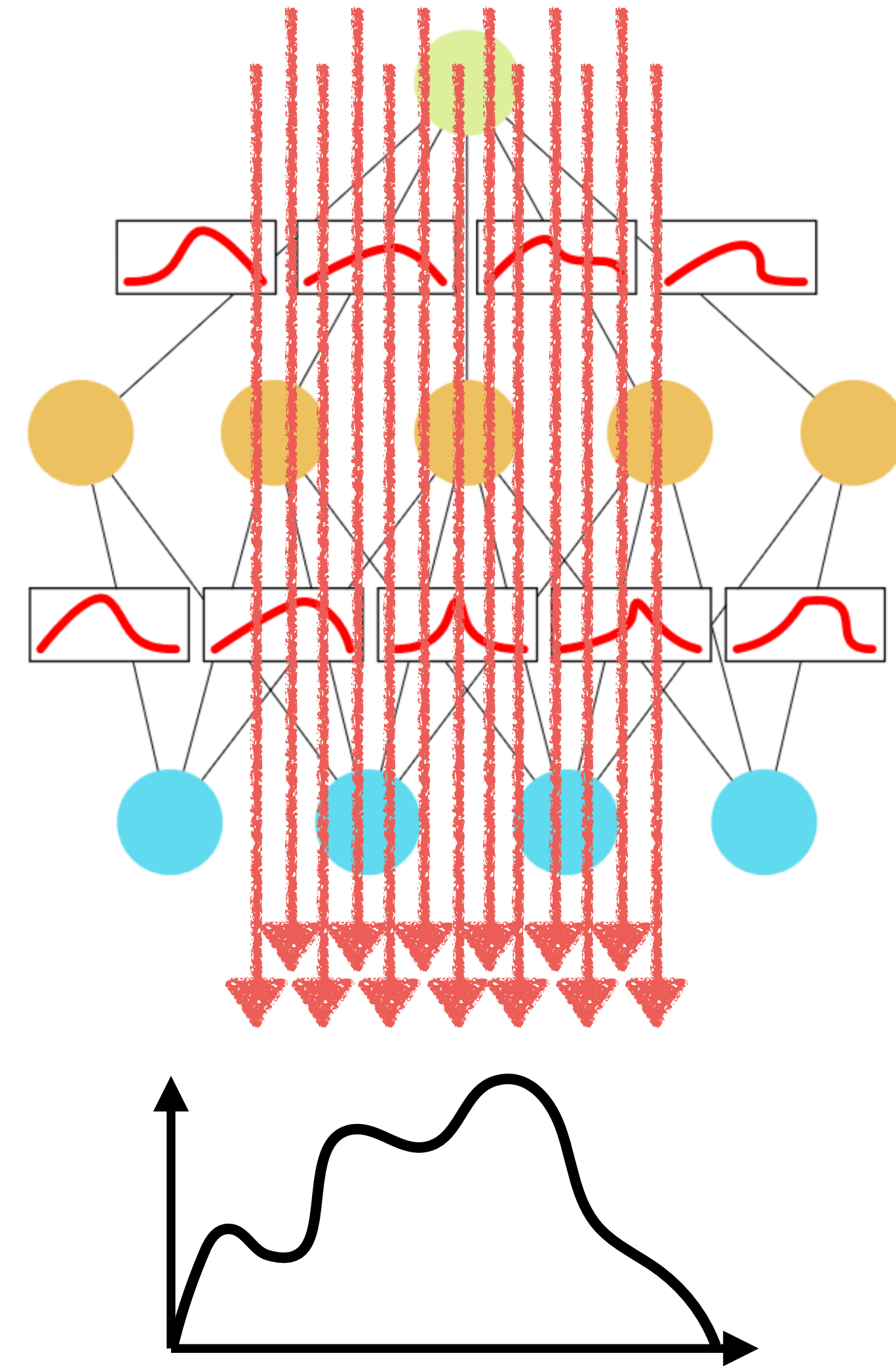
BAYESIAN NEURAL NETWORKS

1. Predict weight distributions with **Bayesian Inference**
2. **Use the BNN** for predictions
 - Use multiple forward passes
 - Each time sample the weights from the weight distributions
 - Analyse the distribution of predictions



BAYESIAN NEURAL NETWORKS

1. Predict weight distributions with **Bayesian Inference**
2. **Use the BNN** for predictions
 - Use multiple forward passes
 - Each time sample the weights from the weight distributions
 - Analyse the distribution of predictions



VARIATIONAL INFERENCE WITH RADICAL APPROXIMATIONS

Deployment on Resource-constrained Embedded Devices

PROBABILISTIC FORWARD PASS

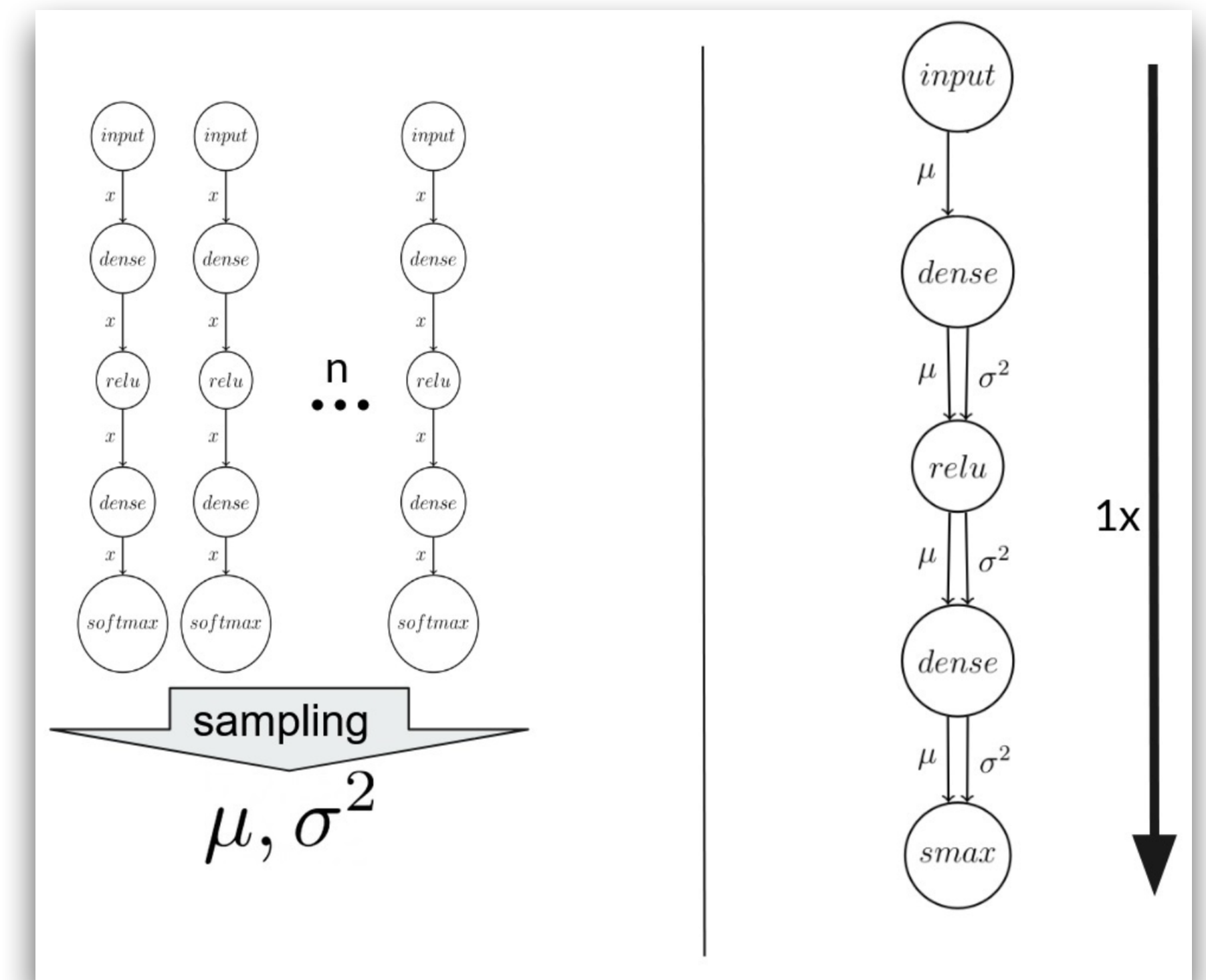
The PFP forces the output activations to be Gaussian distributed.

Reducing the multiple forward passes of VI sampling to a **single**, but more complex, forward pass.

$$E[x] = \mu_x = \frac{\mu_a}{2} \left(1 + \operatorname{erf} \left(\frac{\mu_a}{\sqrt{2\sigma_a^2}} \right) \right) + \sqrt{\frac{\sigma_a}{2\pi}} \exp \left(-\frac{\mu_a^2}{2\sigma_a} \right)$$

$$E[x^2] = \mu_{x^2} = \frac{\sigma_a^2 + \mu_a^2}{2} \left(1 + \operatorname{erf} \left(\frac{\mu_a}{\sqrt{2\sigma_a^2}} \right) \right) + \mu_a \sqrt{\frac{\sigma_a^2}{2\pi}} \exp \left(-\frac{\mu_a^2}{2\sigma_a} \right)$$

PFP ReLU activation function



Sampling-based inference VI (left) vs. Probabilistic Forward Path (Right) [2]

[1] W. Roth, *Variational Inference in Neural Networks using an Approximate Closed-Form Objective*, Workshop on BayesianDL, NIPS 2016.

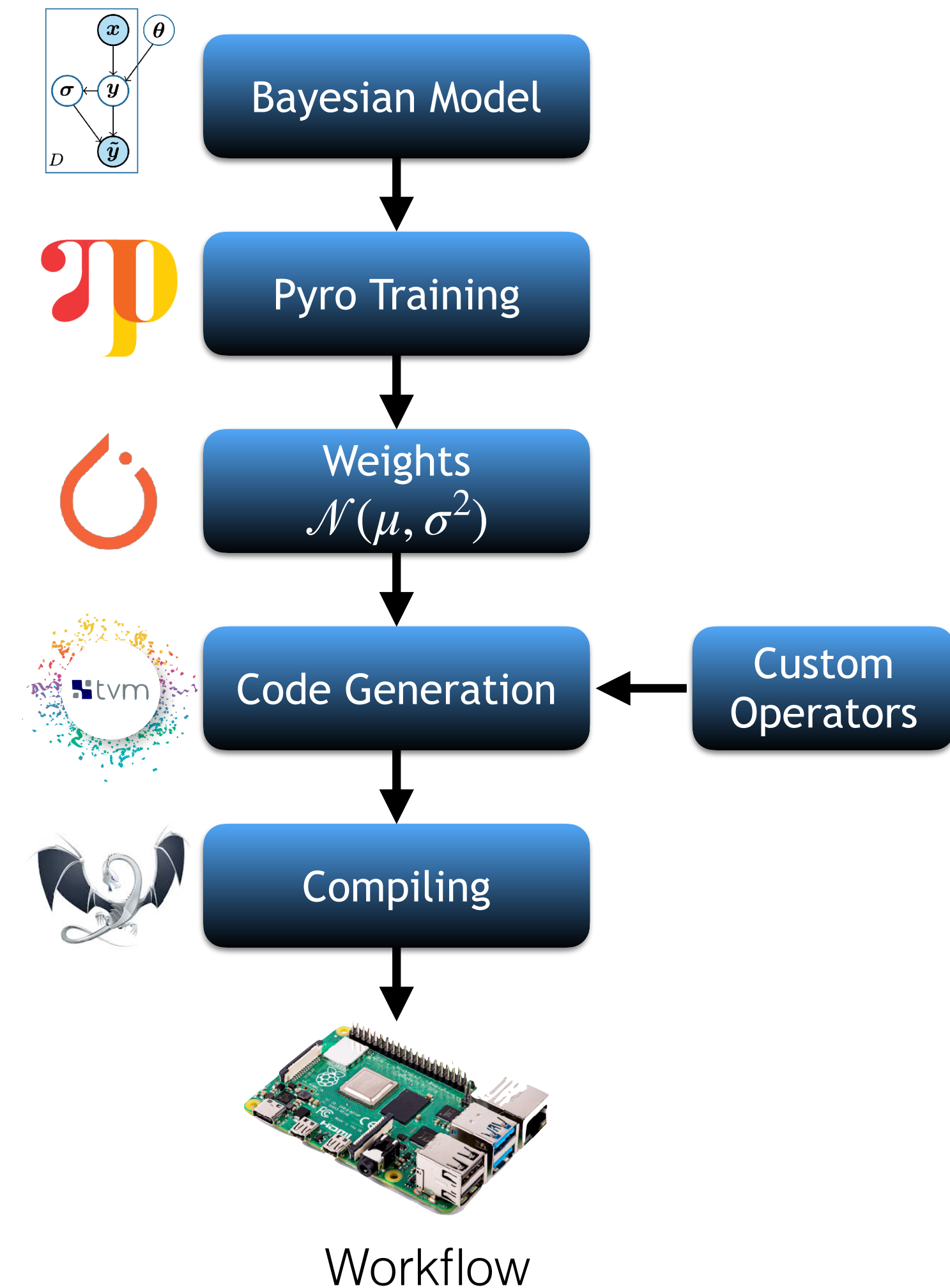
[2] F. Selker, *Optimization of an Approximation Based Approach to Bayesian Neural Networks with TVM on Embedded Hardware*, Master's thesis, Heidelberg University, 2023.

DEPLOYMENT ON EMBEDDED HARDWARE

Exotic operators[1] are not supported by standard libraries

ML Compilers like TVM[2] allow implementation and tuning of custom operators

All operators (convolutions, fully connected layers, activation functions, pool,...) need to be reimplemented and optimized



[1] W. Roth, *Variational Inference in Neural Networks using an Approximate Closed-Form Objective*, Workshop on BayesianDL, NIPS 2016.

[2] <https://github.com/apache/tvm>

PERFORMANCE ON A RASPBERRY PI



Processor	Architecture	Deterministic NN		VI with Pyro	PFP	
		not tuned	tuned		not tuned	tuned
Cortex-A53	MLP	14.0 <i>ms</i>	0.933 <i>ms</i>	-	15.3 <i>ms</i>	4.99 <i>ms</i>
Cortex-A72	MLP	4.59 <i>ms</i>	0.186 <i>ms</i>	738 <i>ms</i>	3.75 <i>ms</i>	0.742 <i>ms</i>
Cortex-A76	MLP	1.64 <i>ms</i>	0.071 <i>ms</i>	309 <i>ms</i>	1.89 <i>ms</i>	0.341 <i>ms</i>
Cortex-A53	LeNet-5	21.1 <i>ms</i>	4.72 <i>ms</i>	-	76.1 <i>ms</i>	37.5 <i>ms</i>
Cortex-A72	LeNet-5	6.89 <i>ms</i>	0.754 <i>ms</i>	1200 <i>ms</i>	21.2 <i>ms</i>	10.1 <i>ms</i>
Cortex-A76	LeNet-5	3.16 <i>ms</i>	0.347 <i>ms</i>	797 <i>ms</i>	9.63 <i>ms</i>	4.02 <i>ms</i>

PERFORMANCE ON A RASPBERRY PI



Processor	Architecture	Deterministic NN		VI with Pyro	PFP	
		not tuned	tuned		not tuned	tuned
Cortex-A53	MLP	14.0 <i>ms</i>	0.933 <i>ms</i>	-	15.3 <i>ms</i>	4.99 <i>ms</i>
Cortex-A72	MLP	4.59 <i>ms</i>	0.186 <i>ms</i>	738 <i>ms</i>	3.75 <i>ms</i>	0.742 <i>ms</i>
Cortex-A76	MLP	1.64 <i>ms</i>	0.071 <i>ms</i>	309 <i>ms</i>	1.89 <i>ms</i>	0.341 <i>ms</i>
Cortex-A53	LeNet-5	21.1 <i>ms</i>	4.72 <i>ms</i>	-	76.1 <i>ms</i>	37.5 <i>ms</i>
Cortex-A72	LeNet-5	6.89 <i>ms</i>	0.754 <i>ms</i>	1200 <i>ms</i>	21.2 <i>ms</i>	10.1 <i>ms</i>
Cortex-A76	LeNet-5	3.16 <i>ms</i>	0.347 <i>ms</i>	797 <i>ms</i>	9.63 <i>ms</i>	4.02 <i>ms</i>

→
AUTO TUNING
10x

→
AUTO TUNING
5x

PERFORMANCE ON A RASPBERRY PI



Processor	Architecture	Deterministic NN		VI with Pyro	PFP	
		not tuned	tuned		not tuned	tuned
Cortex-A53	MLP	14.0 <i>ms</i>	0.933 <i>ms</i>	-	15.3 <i>ms</i>	4.99 <i>ms</i>
Cortex-A72	MLP	4.59 <i>ms</i>	0.186 <i>ms</i>	738 <i>ms</i>	3.75 <i>ms</i>	0.742 <i>ms</i>
Cortex-A76	MLP	1.64 <i>ms</i>	0.071 <i>ms</i>	309 <i>ms</i>	1.89 <i>ms</i>	0.341 <i>ms</i>
Cortex-A53	LeNet-5	21.1 <i>ms</i>	4.72 <i>ms</i>	-	76.1 <i>ms</i>	37.5 <i>ms</i>
Cortex-A72	LeNet-5	6.89 <i>ms</i>	0.754 <i>ms</i>	1200 <i>ms</i>	21.2 <i>ms</i>	10.1 <i>ms</i>
Cortex-A76	LeNet-5	3.16 <i>ms</i>	0.347 <i>ms</i>	797 <i>ms</i>	9.63 <i>ms</i>	4.02 <i>ms</i>

→
AUTO TUNING
10X

→
AUTO TUNING
5X

→
OVERHEAD PFP 5X/10X

PERFORMANCE ON A RASPBERRY PI



Processor	Architecture	Deterministic NN		VI with Pyro	PFP	
		not tuned	tuned		not tuned	tuned
Cortex-A53	MLP	14.0 ms	0.933 ms	-	COST REDUCTION 	
Cortex-A72	MLP	4.59 ms	0.186 ms	738 ms		
Cortex-A76	MLP	1.64 ms	0.071 ms	309 ms		
Cortex-A53	LeNet-5	21.1 ms	4.72 ms	-	76.1 ms	37.5 ms
Cortex-A72	LeNet-5	6.89 ms	0.754 ms	1200 ms	21.2 ms	10.1 ms
Cortex-A76	LeNet-5	3.16 ms	0.347 ms	797 ms	100x ms	4.02 ms


AUTO TUNING
10x


AUTO TUNING
5x


OVERHEAD PFP 5x/10x

SUMMARY

Automatic Compression (Galen)

Guide search with

Sensitivity Information

Latency measurements

Compression methods can balance each other.

Bayesian Neural Networks

Report probabilities and can say “I don’t know”.

Computational cost is significantly higher than their classical counterparts.

Approximations can be used to reduce this cost drastically.

Deep Learning Compilers (TVM)

Custom Operators

Auto-tuning

