



Energy-Constrained Compression for Deep Neural Networks via Weighted Sparse Projection and Layer Input Masking



Haichuan Yang¹, Yuhao Zhu¹, Ji Liu^{2,1}

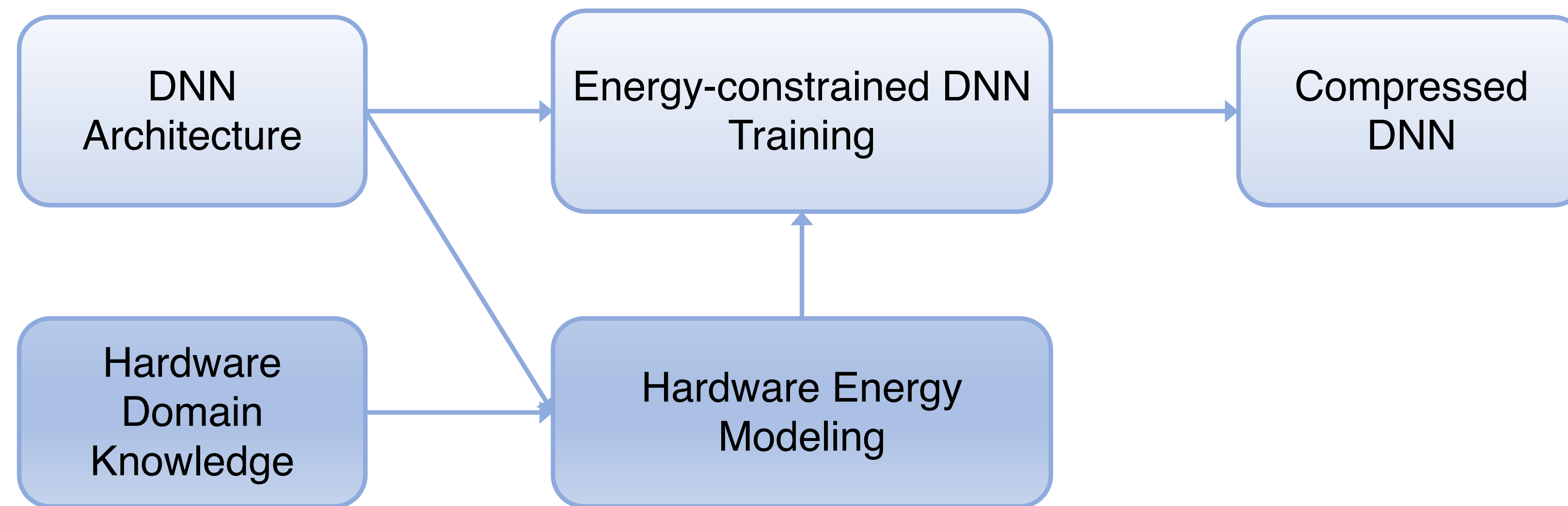
¹University of Rochester, ²Kwai AI Lab at Seattle

Motivation

- Deep Neural Networks (DNNs) are increasingly deployed in highly energy-constrained environments.
- Compression methods like pruning or quantization are proposed to reduce the redundant parameters.
- Existing work: barely consider the hardware and energy constraint in pruning.
- This work: use the hardware energy model to guide the pruning.

Overview

Overview of the framework:



The overall objective:

$$\begin{aligned}
 & \min_W \mathcal{L}(W) \\
 & \text{s.t. } E(W) \leq E_{\text{budget}}
 \end{aligned}$$

DNN Parameters ← W DNN Training Loss

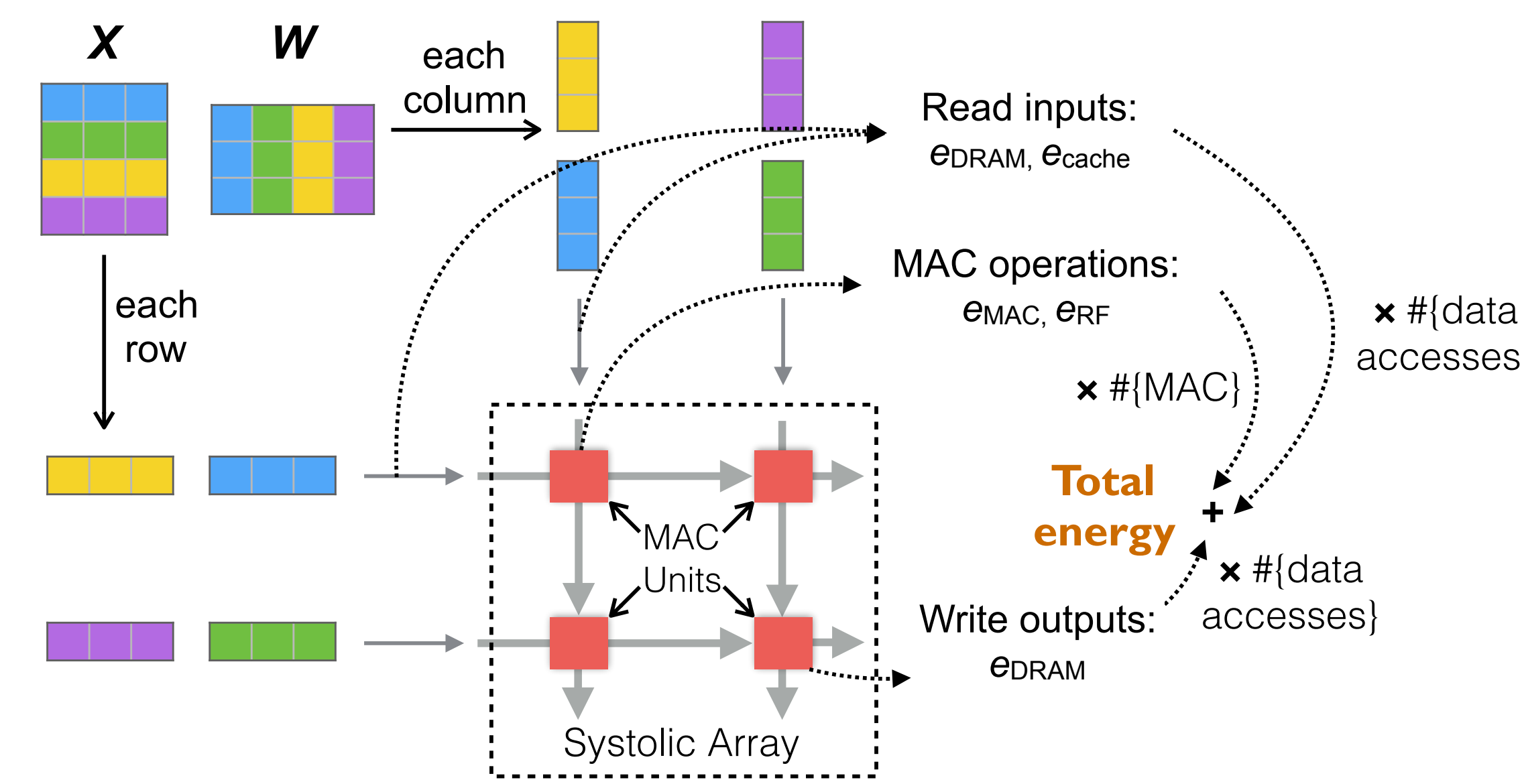
↓
Energy Modeling

Modeling DNN Inference Energy Consumption

Targeted Hardware: systolic-array-based.

Count the hardware operations:

- Data reuse in parallel MACs;
- Zero operands are skipped.



Energy Model: piece-wise linear function $E(W) =$

$$\sum_{i=1}^L \underbrace{\alpha^{(i)}}_{\text{unit energy cost (in-cache)}} \min(\underbrace{k}_{\text{cache size}}, \|W^{(i)}\|_0) + \underbrace{\beta^{(i)}}_{\text{unit energy cost (out-cache)}} \max(0, \|W^{(i)}\|_0 - k) + \underbrace{\gamma^{(i)}}_{\text{unit energy cost (cache-independent)}} \|W^{(i)}\|_0 + \underbrace{\delta}_{\text{constant overhead}}$$

Energy-Constrained DNN Training

Basic Idea: Projected-SGD

$$W \leftarrow \text{P} (W - \eta \nabla \mathcal{L}(W))$$

SGD

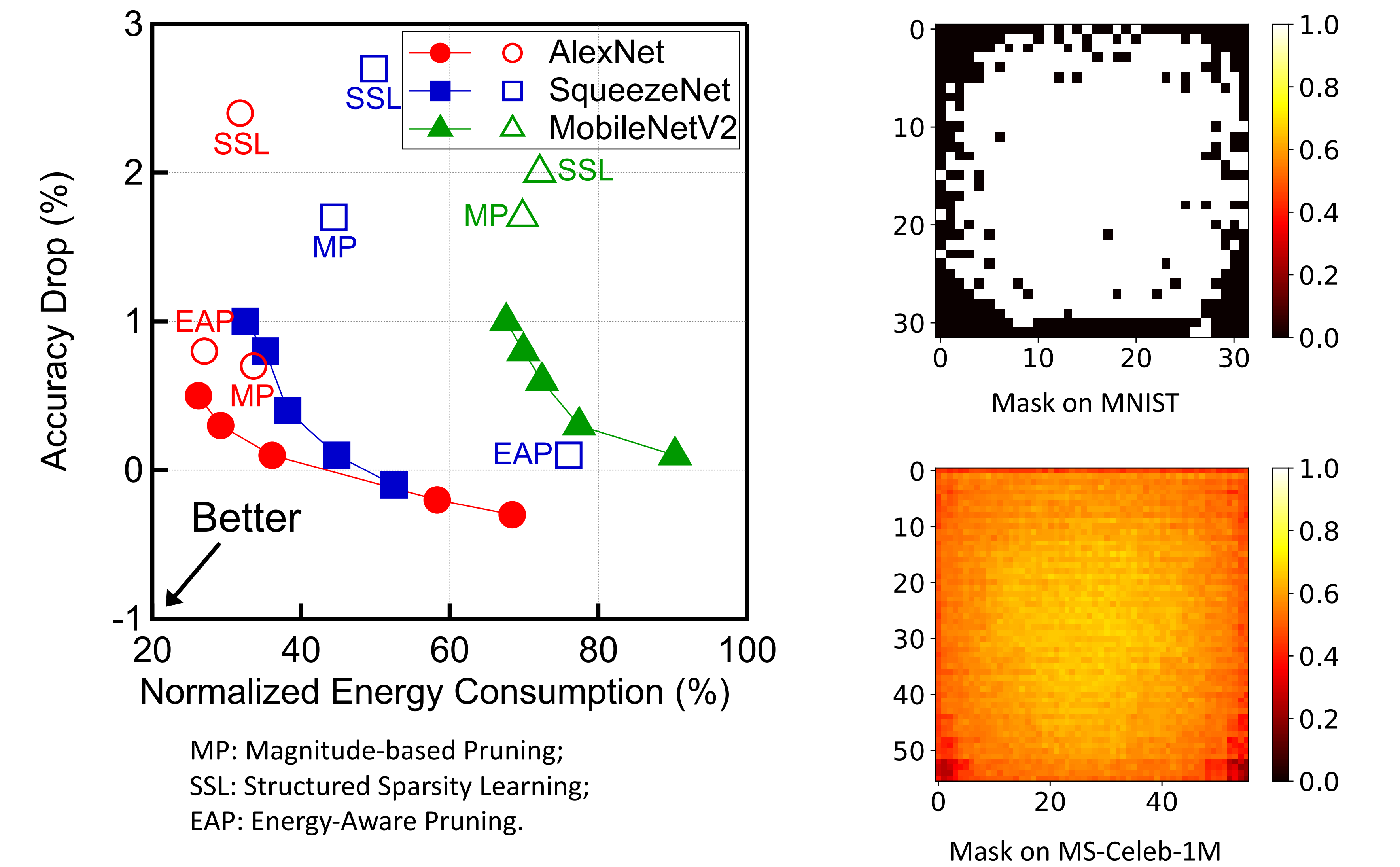
Projection:
 $\text{P}(Z) := \underset{W: E(W) \leq E_{\text{budget}}}{\text{argmin}} \|W - Z\|$

Theorem 1
 $\iff$

0/1 knapsack

Additional trick: design a trainable binary mask to increase the sparsity of the input mask.

Experiment Results



DNNs	AlexNet 26%					SqueezeNet 38%				MobileNetV2 68%		
Energy Budget	MP	SSL	EAP	NetAdapt	Ours	MP	SSL	EAP	Ours	MP	SSL	Ours
Methods	0.7%	2.4%	0.8%	4.4%	0.5%	1.7%	2.7%	0.1%	0.4%	1.7%	2.0%	1.0%
Accuracy Drop	34%	32%	27%	26%	26%	44%	50%	76%	38%	70%	72%	68%
Nonzero Weights Ratio	8%	35%	9%	10%	31%	34%	61%	28%	48%	52%	63%	63%

DNNs@Dataset	LeNet-5@MNIST 17%						MobileNetV2@MS-Celeb-1M 60%		
Energy Budget	MP	SSL	NetAdapt	SBP	BC	Ours	MP	SSL	Ours
Methods	1.5%	1.5%	0.6%	1.5%	2.2%	0.5%	1.1%	0.7%	0.2%
Accuracy Drop	18%	20%	18%	22%	26%	17%	66%	72%	60%
Energy									

SBP: Structured Bayesian Pruning; BC: Bayesian Compression

Conclusion

- Propose an end-to-end energy-aware model compression framework;
- Outperform state-of-the-arts;
- First work to involve the hardware model in model compression.