

FLightNNs: Lightweight Quantized Deep Neural Networks for Fast and Accurate Inference

Ruizhou Ding, Zeye Liu, Ting-Wu Chin, Diana Marculescu, and R. D. (Shawn) Blanton
{rding,zeyel,tingwuc,dianam,rblanton}@andrew.cmu.edu
Carnegie Mellon University, Pittsburgh, U.S.A.

ABSTRACT

To improve the throughput and energy efficiency of Deep Neural Networks (DNNs) on customized hardware, lightweight neural networks constrain the weights of DNNs to be a limited combination (denoted as $k \in \{1, 2\}$) of powers of 2. In such networks, the multiply-accumulate operation can be replaced with a single shift operation, or two shifts and an add operation. To provide even more design flexibility, the k for each convolutional filter can be optimally chosen instead of being fixed for every filter. In this paper, we formulate the selection of k to be differentiable, and describe model training for determining k -based weights on a per-filter basis. Over 46 FPGA-design experiments involving eight configurations and four data sets reveal that lightweight neural networks with a flexible k value (dubbed FLightNNs) fully utilize the hardware resources on Field Programmable Gate Arrays (FPGAs), our experimental results show that FLightNNs can achieve $2\times$ speedup when compared to lightweight NNs with $k = 2$, with only 0.1% accuracy degradation. Compared to a 4-bit fixed-point quantization, FLightNNs achieve higher accuracy and up to $2\times$ inference speedup, due to their lightweight shift operations. In addition, our experiments also demonstrate that FLightNNs can achieve higher computational energy efficiency for ASIC implementation.

CCS CONCEPTS

• Computing methodologies → Machine learning; • Hardware → Electronic design automation.

ACM Reference Format:

Ruizhou Ding, Zeye Liu, Ting-Wu Chin, Diana Marculescu, and R. D. (Shawn) Blanton. 2019. FLightNNs: Lightweight Quantized Deep Neural Networks for Fast and Accurate Inference. In *The 56th Annual Design Automation Conference 2019 (DAC '19)*, June 2–6, 2019, Las Vegas, NV, USA. ACM, New York, NY, USA, 6 pages. <https://doi.org/10.1145/3316781.3317828>

1 INTRODUCTION

Emerging vision, speech and natural language applications have widely adopted deep learning models and, as a result, have achieved state-of-the-art accuracy. Furthermore, recent industrial efforts have focused on implementing the models on mobile devices [1]. However, real-time applications based on these deep models may incur unacceptably large latencies and can easily drain the battery on energy-limited devices. For example, smartphones can only run

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

DAC '19, June 2–6, 2019, Las Vegas, NV, USA

© 2019 Copyright held by the owner/author(s). Publication rights licensed to ACM.

ACM ISBN 978-1-4503-6725-7/19/06...\$15.00

<https://doi.org/10.1145/3316781.3317828>

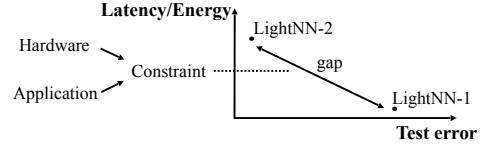


Figure 1: A discrete Pareto-optimal curve for LightNN models w.r.t. test error and latency/energy. More continuous Pareto-optimal points are needed to adapt to the latency/energy constraints determined by the hardware and application.

the AlexNet-based object detection for one hour [2]. Therefore, prior research has proposed model compression techniques including pruning and quantization to satisfy the stringent energy and speed requirements [3].

One of the recently proposed quantization approaches, LightNN, constrains the weights of DNNs to be a sum of k powers of 2, and therefore can use shift and add operations to replace the multiplications between activations and weights [4]. For LightNN-1¹, all the multiplications of the DNNs will be replaced by a shift operation, while for LightNN-2, two shifts and an add replace the multiplication. Since shift operations are much more lightweight on customized hardware (e.g., FPGA or ASIC), LightNNs can achieve faster speed and lower energy consumption, and generally maintain accuracy for over-parameterized models [4, 5]. Although LightNNs provide better energy-efficiency, they lack the flexibility to provide fine-grained trade-offs between energy and accuracy. As shown in Fig. 1, the energy efficiency for these models also exhibits gaps, making the Pareto front of accuracy and energy discrete. However, a continuous accuracy and energy/latency trade-off is an important feature for designers to target different market segments (e.g., IoT devices, edge devices, and mobile devices).

To provide a more flexible Pareto front for the LightNN framework, we propose to equip each convolutional filter with the freedom to use a different number of shift-and-add operations to approximate multiplications. Specifically, we introduce a set of free variables $\mathbf{k} = \{k_1, \dots, k_F\}$ where each element represents the number of shift-and-add for the corresponding convolutional filter. As a result, a more contiguous Pareto front can be achieved. For example, if we constrain $\mathbf{k} \in \{1, 2\}^F$, then the throughput and energy consumption of the new model will sit between LightNN-1 ($\mathbf{k} = \{1\}^F$) and LightNN-2 ($\mathbf{k} = \{2\}^F$). Formally, we are solving $\min_{\mathbf{w}, \mathbf{k}} \mathcal{L}(\mathbf{w}, \mathbf{k})$, where $\mathcal{L}$ is the loss function and $\mathbf{w}$ is the weights vector. However, the commonly adopted stochastic gradient descent (SGD) algorithm does not apply in this case since $\mathcal{L}$ is non-differentiable w.r.t. $\mathbf{k}$. In this paper, we propose a *differentiable training algorithm* which enables end-to-end optimization with standard SGD. The resulting network is dubbed *FLightNN* for its flexible $\mathbf{k}$ values.

¹LightNN- k quantizes weights to be the sum of k powers of 2.

2 RELATED WORK

Prior work has extensively explored approaches to reduce latency and energy consumption of DNNs on hardware, through both algorithmic [2, 6] and hardware [7, 8] efforts. Since the latency and energy consumption of DNNs generally stem from computational cost and memory accesses, prior work in the algorithmic domain mainly focuses on the reduction of FLOPs and model size. Some work reduces the number of parameters through weight pruning [9], while some other work introduces structural sparsity via filter pruning for Convolutional Neural Networks (CNNs) [10] to enable speedup on general hardware platforms incorporating CPUs and GPUs. To reduce the model size, previous work has also conducted neural architecture search with energy constraint [2]. In addition to algorithmic advances, prior art has also proposed methodologies to achieve fast and energy-efficient DNNs. Some previous work proposes the co-design of the hardware platform and the architecture of the neural network running on it [11]. Some work proposes more lightweight DNN units for faster inference on general-purpose hardware [12], while others propose hardware-friendly DNN computation units to enable energy-efficient implementation on customized hardware [13].

By reducing the weight and activation precision, DNN quantization has proved to be an effective technique to improve the speed and energy efficiency of DNNs on customized hardware, due to its lower computational cost and fewer memory accesses [14]. Gupta *et al.* show that a DNN with 16-bit fixed-point representation can achieve competitive accuracy compared to the full-precision network [14]. In the same vein, Zhou *et al.* explored the DNN accuracy *w.r.t.* a wide range of bit widths [15]. These uniform quantization approaches enable fixed-point hardware implementation for DNNs. Courbariaux *et al.* propose BinaryConnect, which uses only 1 bit for the DNN parameters, turning multiplications into XNOR operations on customized hardware [16]. However, these models require an over-parameterized model size to maintain a high accuracy [4].

LightNNs constrain the model weights to be a power of 2, or the sum of a limited number of powers of 2 [4], while the activations use fixed-point quantization. Therefore, the multiplication between weights and activations can be implemented in hardware by shift operations and fixed-point additions. Compared to DNNs with fixed-point quantization, LightNNs replace the fixed-point multipliers by more lightweight shift operators, or shift and additions. Since the shift operators can be implemented using Look-Up Table (LUT) on FPGA while fixed-point multipliers require Digital Signal Processing (DSP) units, LightNNs can have higher inference speed than fixed-point DNNs when run on DSP-bounded FPGAs. In addition, in an ASIC implementation, shift operations are more lightweight than multiplications, making LightNNs more energy and area efficient than fixed-point DNNs.

However, LightNNs use a single k value (*i.e.*, the number of shifts per multiplication) across the whole network, and therefore lack flexibility to provide a fine-grained energy/latency and accuracy trade-off for hardware designers. Therefore, we propose FLIGHTNNs which use customized k values for each convolutional filter to enable a more continuous Pareto front. Recent work has explored the idea of differentiable training for architecture search [17] and neural network pruning [18]. In this paper, we propose an end-to-end differentiable training algorithm for FLIGHTNNs via approximate

gradient computation for non-differentiable operations and regularization to encourage sparsity. Moreover, the proposed differentiable training approach uses gradual quantization, which can achieve higher accuracy than LightNN-1 without increasing latency. In summary, this paper has the following key contributions:

(i) We propose a differentiable training algorithm for FLIGHTNNs, which provides a continuous Pareto front for hardware designers to search for a highly accurate model under the hardware resource constraints.

(ii) The differentiable training for FLIGHTNNs enables gradual quantization, and further pushes forward the Pareto-optimal curve.

3 LIGHTNN OVERVIEW

As a quantized DNN model, LightNNs constrain the weights of a network to be the sum of k powers of 2, denoted as LightNN- k . Thus, the multiplications between weights and activations can be implemented with k shift operations and $k-1$ additions. Specifically, LightNN-1 constrains the weights to be a power of 2, and only uses a shift for a multiplication. The approximation function used by LightNN- k to quantize a full-precision weight w can be formulated in a recursive way: $Q_k(w) = Q_{k-1}(w) + Q_1(w - Q_{k-1}(w))$ for $k > 1$, where $Q_1(w) = \text{sign}(w) \times 2^{\lfloor \log(|w|) \rfloor}$ which rounds the weight w to a nearest power of 2.

LightNNs are trained with a modified backpropagation algorithm. In the forward phase of each training iteration, the parameters are first approximated using the Q_k function. Then, in the backward phase, the gradients of loss *w.r.t.* quantized weights are computed, and applied to the full-precision weights in the weight update phase. LightNNs have been proved to be accurate and energy-efficient on customized hardware [4]. LightNN-2 can generally have an accuracy close to full-precision DNNs, while LightNN-1 can achieve higher energy efficiency than LightNN-2. Due to the nature of the discrete k values, there exists a gap between LightNN-1 and LightNN-2 *w.r.t.* accuracy and energy. We propose to customize the k values for each convolutional filter, and thus, achieve a smoother energy-accuracy trade-off to provide hardware designers with more design options.

4 DIFFERENTIABLE TRAINING FOR FLIGHTNNs

In this section, we first define the quantization function, and then introduce the end-to-end training algorithm for FLIGHTNNs, equipped with a regularization loss to penalize large k values.

4.1 Quantization function

We first denote the i^{th} filter of the network as $\mathbf{w}_i$ and the quantization function for the filter $\mathbf{w}_i$ as $Q_k(\mathbf{w}_i|\mathbf{t})$, where $k = \max_j k_j$ is the maximum number of shifts used for this network, and vector $\mathbf{t}$ is a latent variable that controls the approximation (*e.g.*, some threshold value). Also, we denote the residual resulting from the approximation as $\mathbf{r}_{i,k} = \mathbf{w}_i - Q_k(\mathbf{w}_i|\mathbf{t})$. Then, we formally define the quantization function as follows:

$$Q_k(\mathbf{w}_i|\mathbf{t}) = \begin{cases} 0, & \text{if } k = 0 \\ \sum_{j=0}^{k-1} \mathbb{1}(\|\mathbf{r}_{i,j}\|_2 > t_j) R(\mathbf{r}_{i,j}), & \text{if } k \geq 1 \end{cases}$$

where $R(x) = \text{sign}(x) \times 2^{\lfloor \log(|x|) \rfloor}$ rounds the input variable to a nearest power of 2, and $\lfloor \cdot \rfloor$ is a rounding-to-integer function. This quantization flow is shown in Fig. 2. To interpret the thresholds t , t_0 determines whether this filter is pruned out, and t_1 determines

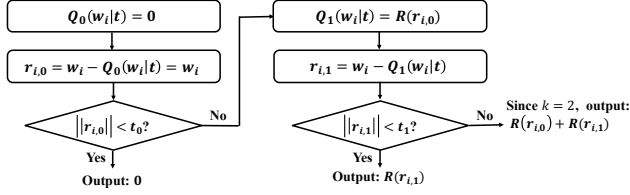


Figure 2: Quantization flow for $k = 2$.

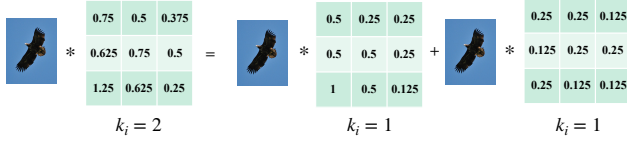


Figure 3: Equivalent conversion from a convolution with a $k_i > 1$ filter to k_i convolutions each with a $k_i = 1$ filter. This transforms the hardware implementation of the FLIGHTNN into LightNN-1.

whether one shift is enough, *etc.* Then, the number of shifts for the i -th filter is $k_i = \sum_{j=0}^{k-1} \mathbb{1}(\|r_{i,j}\|_2 > t_j)$. Therefore, choosing k_i per filter is equivalent to finding optimal thresholds t .

The FLIGHTNN quantization approach targets efficient hardware implementation. Instead of assigning a customized k_i for each weight, FLIGHTNNs have customized k_i values per filter, and therefore preserve the structural sparsity. As shown in Fig. 3, the convolution with a $k_i = 2$ filter can be equivalently converted to the sum of two convolutions each with a $k_i = 1$ filter. Thus, FLIGHTNNs can be efficiently implemented as LightNN-1 with an extra summation of feature maps per layer.

4.2 Differentiable training

Instead of picking the thresholds t by hand, we consider them as trainable parameters. Therefore, the loss function $\mathcal{L}(\mathbf{w}^2, \mathbf{t})$ is a function of both weights and thresholds. Similar to prior work on DNN quantization [15, 16], we use the straight-through estimator (STE) [19] to compute $\frac{\partial \mathcal{L}}{\partial \mathbf{w}_i^q}$. By defining $\frac{\partial \mathbf{w}_i^q}{\partial \mathbf{w}_i} = 1$ where $\mathbf{w}_i^q = Q_k(\mathbf{w}_i|\mathbf{t})$ is the quantized $\mathbf{w}_i$; therefore, we have $\frac{\partial \mathcal{L}}{\partial \mathbf{w}_i} = \frac{\partial \mathcal{L}}{\partial \mathbf{w}_i^q}$. $\frac{\partial \mathbf{w}_i^q}{\partial \mathbf{w}_i} = \frac{\partial \mathcal{L}}{\partial \mathbf{w}_i^q}$, which becomes a differentiable expression.

To compute the gradient for thresholds, *i.e.*, $\frac{\partial \mathbf{w}_i^q}{\partial t_j}$, we relax the indicator function $g(x, t_j) = \mathbb{1}(x > t_j)$ to a sigmoid function [20], $\sigma(\cdot)$, when computing gradients, *i.e.*, $\hat{g}(x, t_j) = \sigma(x - t_j)$. In addition, we use STE to compute the gradient for $R(x)$. Thus, the gradient $\frac{\partial \mathbf{w}_i^q}{\partial t_j}$ can be computed by:

$$\begin{aligned} \frac{\partial Q_{k_i}(\mathbf{w}_i|\mathbf{t})}{\partial t_j} &= \sum_{l=0}^{k_i-1} \frac{\partial \sigma(\|r_{i,l}\|_2 - t_l)}{\partial t_j} R(r_{i,l}) + \sigma(\|r_{i,l}\|_2 - t_l) \frac{\partial R(r_{i,l})}{\partial t_j} \\ &= \sum_{l=0}^{k_i-1} \sigma'(\|r_{i,l}\|_2 - t_l) \left(\frac{\partial \|r_{i,l}\|_2}{\partial t_j} - \frac{\partial t_l}{\partial t_j} \right) R(r_{i,l}) + \sigma(\|r_{i,l}\|_2 - t_l) \frac{\partial r_{i,l}}{\partial t_j} \end{aligned}$$

where $\frac{\partial \|r_{i,l}\|_2}{\partial t_j}$ and $\frac{\partial r_{i,l}}{\partial t_j}$ are 0 for $l < j$; otherwise, they can be computed with the result of $\frac{\partial Q_l(\mathbf{w}_i|\mathbf{t})}{\partial t_j}$. $\frac{\partial t_l}{\partial t_j} = \mathbb{1}(l = j)$.

²The bias term is omitted for simplicity.

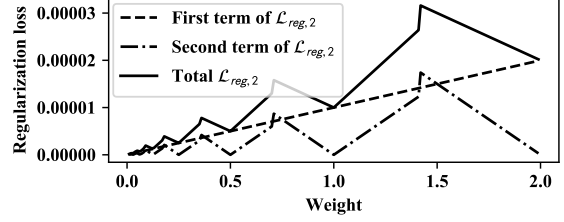


Figure 4: Regularization loss curve w.r.t. weight value.

Algorithm 1: FLIGHTNN Training Epoch

Input: Training dataset $(\mathbf{x}, \mathbf{y})$, where $\mathbf{x}$ is input and $\mathbf{y}$ is label; parameters after the $(p-1)$ -th iteration: $\mathbf{w}_{p-1}$ (weights), $\mathbf{b}_{p-1}$ (biases), and quantization thresholds $\mathbf{t}_{p-1}$; quantization function $Q_k(\mathbf{w}|\mathbf{t})$; DNN forward computation function $g(\mathbf{x}, \mathbf{w}, \mathbf{b})$; maximum k value used for all filters; regularization loss coefficients λ ; learning rate η .

Output: Updated weights $\mathbf{w}_p$, biases $\mathbf{b}_p$ and thresholds $\mathbf{t}_p$.

for each mini-batch of $\mathbf{x}$, $\mathbf{y}$ do

1. **Quantize weights:** $\mathbf{w}^q = Q_k(\mathbf{w}_{p-1}|\mathbf{t}_{p-1})$
2. **Forward:** compute intermediate results and cross entropy loss function $\mathcal{L}_{CE}$ with $g(\cdot)$, $\mathbf{w}^q$, $\mathbf{b}_{p-1}$, and mini-batch of $\mathbf{x}$; compute regularization loss $\mathcal{L}_{reg,k}$ with λ and $\mathbf{w}_{p-1}$; get the total loss $\mathcal{L}_{total} = \mathcal{L}_{CE} + \mathcal{L}_{reg,k}$
3. **Backward:** compute derivatives $\frac{\partial \mathcal{L}_{total}}{\partial \mathbf{w}^q}$, $\frac{\partial \mathcal{L}_{total}}{\partial \mathbf{b}_{p-1}}$, and $\frac{\partial \mathcal{L}_{total}}{\partial \mathbf{t}_{p-1}}$
4. **Update parameters:** $\mathbf{w}_p = \mathbf{w}_{p-1} - \eta \frac{\partial \mathcal{L}_{total}}{\partial \mathbf{w}^q}$; $\mathbf{b}_p = \mathbf{b}_{p-1} - \eta \frac{\partial \mathcal{L}_{total}}{\partial \mathbf{b}_{p-1}}$; $\mathbf{t}_p = \mathbf{t}_{p-1} - \eta \frac{\partial \mathcal{L}_{total}}{\partial \mathbf{t}_{p-1}}$

end

4.3 Regularization

To encourage smaller k_i for the filters, we also add a regularization loss: $\mathcal{L}_{reg,k}(\mathbf{w}) = \sum_{j=0}^{k-1} \lambda_j \sum_i \|r_{i,j}\|_2$ where λ_j performs as a handle to balance accuracy and model sparsity. This regularization loss is the sum of several group Lasso losses, since they can introduce structural sparsity [10]. The first item $\lambda_0 \sum_i \|r_{i,0}\|_2 = \lambda_0 \sum_i \|\mathbf{w}_i\|_2$ is used to prune the whole filters out, while the other items ($j > 0$) regularize the residuals. Fig. 4 shows the two items of regularization loss and their sum for the case $k = 2$, with $\lambda_0 = 1e-5$ and $\lambda_1 = 3e-5$. Therefore, the total loss for training a FLIGHTNN is: $\mathcal{L}_{total}(\mathbf{w}, \mathbf{t}) = \mathcal{L}_{CE}(\mathbf{w}, \mathbf{t}) + \mathcal{L}_{reg,k}(\mathbf{w})$.

The new training algorithm is summarized in Algo. 1. This is the same as the conventional backpropagation algorithm for full-precision DNNs, except that in the forward phase, the weights are quantized given the thresholds $\mathbf{t}$. Then, due to the differentiability of the quantization function w.r.t. $\mathbf{w}$ and $\mathbf{t}$, one can compute their gradients and update their values in each training iteration.

5 EXPERIMENTAL RESULTS

In this section, we first introduce the experiment setup. Then, we show the accuracy results of different quantized DNN models by software training, as well as their throughput on the FPGA and energy efficiency on the ASIC, to verify the effectiveness of FLIGHTNNs.

Table 1: Network settings. “Depth” is the number of convolutional layers in the network. “Width” is the number of convolutional filters of the largest layer.

Network ID	Parameters	Structure	Depth	Width
1	0.08M	VGG	7	64
2	0.7M	ResNet	18	128
3	4.6M	VGG	7	512
4	0.03M	VGG	4	64
5	0.1M	VGG	4	128
6	0.7M	ResNet	18	128
7	2.8M	ResNet	18	256
8	1.8M	ResNet	10	256

5.1 Setup

We conduct experiments on both small and large CNNs for CIFAR-10, SVHN, CIFAR-100 and ImageNet datasets. The eight adopted network configurations are shown in Table 1. To explore the FLightNN performance on different types of network structures, we use a VGG structure with a series of stacked convolutional layers for Network 1, 3, 4 and 5, and adopt the ResNet structure with skip connections across layers for network 2, 6, 7 and 8. Networks 1, 2 and 3 are used for experiments on CIFAR-10; networks 4 and 5 are used for SVHN; networks 6 and 7 are used for CIFAR-100; the last one, network 8, is used for ImageNet. For all networks, each convolutional layer is followed by a batch normalization layer and a Leaky ReLU activation function [21], and optionally followed by a max-pooling layer. We use the Adam optimizer [22] to train the network. For each of the networks, we train different quantized models including full-precision DNNs, fixed-point DNNs with 4-bit weights and 8-bit activations, LightNN-2 with 8-bit weights and 8-bit activations, LightNN-1 with 4-bit weights and 8-bit activations, and FLightNNs with 8-bit activations. Due to large training times and limitations in computing resources, we train the ImageNet dataset on a ResNet-10 with reduced width (*i.e.*, network 8), for LightNN-1, LightNN-2 and FLightNNs. For all FLightNNs, we initialize the thresholds t to 0, and set the largest shifts k as 2. For all, except the 32-bit full-precision model, we use 8-bit fixed-point quantization for the activations. By varying λ , we can have different accuracy-throughput or accuracy-energy trade-offs for FLightNNs. All these networks are trained in software through PyTorch.

5.2 Accuracy-throughput trade-off on FPGA

To show the accuracy-throughput trade-off of the models, we implement the inference of each network’s largest convolutional layer for each of the quantized DNN models on FPGA since prior work has shown that convolution operations typically take over 90% of the computation time of a CNN [23]. Our implementation is built on the Xilinx Zynq ZC706 evaluation board. Its working frequency is 100 MHz. Pre-synthesis is executed on an Intel i7-4790 CPU (3.6GHz) with 16GB RAM. We use Vivado HLS [24] for FPGA implementation. The C code of DNN designs are parallelized by adding HLS-defined pragma and the parallel version is validated with the Vivado HLS timing analysis tool. To make a fair comparison, the same pragma and directives are used for full-precision, fixed-point DNNs, LightNNs and FLightNNs, and we follow the same scheduling settings as prior work [5]. Batched inference is

adopted, and the maximum batch size without running out of FPGA resources is set to obtain the highest throughput.

Tables 2, 3, 4 and 5 show the accuracy and throughput comparison for full-precision DNNs, fixed-point DNNs, LightNNs and FLightNNs. For all the experimented datasets, LightNNs show the advantage of flexible accuracy-speed trade-offs. In most of the networks (*e.g.*, networks 1, 3, 6 and 7), FLightNNs can achieve an accuracy close to LightNN-2, but have much higher speedup than LightNN-2. Thus, FLightNNs provide continuous trade-offs for accuracy and speed. Compared to the fixed-point quantization, FLightNNs can achieve higher accuracy, and up to 2.0 $\times$, 1.8 $\times$ and 1.8 $\times$ speedup for CIFAR-10, SVHN, CIFAR-100 datasets, respectively. This is because the multiplication is replaced by shift operators, which require only LUT resources on FPGA while the multipliers require DSP units which are generally more scarce than LUT. Therefore, the computation for FLightNNs allows larger batch sizes than that of fixed-point DNNs, increasing data parallelism, and thus, improving the throughput.

It is also interesting to note that by comparing some FLightNNs (*e.g.*, FL_{1a}, FL_{2a}, FL_{3a}, FL_{6a} and FL_{7a}) with LightNN-1, we find that FLightNNs can achieve higher accuracy with the same or even lower storage as LightNN-1. This is because initially FLightNNs quantize all the filters with two shifts (since t is initialized as 0), and gradually add constraints to the filters. This gradual quantization may be better than training a network with only one shift from scratch, as LightNN-1 does. The benefit of gradual quantization has also been observed by prior work [25] which shows that gradually imposing quantization constraints can achieve better accuracy than directly quantizing with a strict constraint.

Table 6 shows the FPGA resources utilization for networks 7 and 8. Since full-precision and fixed-point DNNs require DSP for both multiplication and addition, while LightNNs and FLightNNs only need DSP for addition, full-precision and fixed-point DNNs have larger DSP resource utilization. Compared to full-precision DNNs which use 32-bit floating point operations, fixed-point DNNs only use 4-bit weights and 8-bit activations, and therefore consume fewer DSP units. LightNNs and FLightNNs use LUT to implement the multipliers, and have a higher utilization of LUT than full-precision and fixed-point DNNs. However, the performance of (F)LightNNs is not bounded by LUT resources since the maximum usage of LUT by LightNN-2 is only 42% and 17% for networks 7 and 8, respectively. Instead, the memory resource (BRAM) bounds the performance for (F)LightNNs, while for full-precision and fixed-point DNNs, the performance is bounded by both BRAM and DSP.

5.3 Accuracy-energy trade-off on ASIC

For all quantized DNNs, we designed pipelined implementations with one stage per neuron, where the computation unit is reused for each neuron. A 65nm commercial standard library is adopted. The Synopsys Design Compiler [26] is used to generate the gate-level netlist of the computation units. The power consumption of all computation operations within one layer is calculated using Synopsys Primetime. We keep all the DNN architectures implemented in an unoptimized fashion because our main objective is to compare how different quantized DNNs impact computational energy.

The accuracy and computational energy trade-offs for the quantized DNN models are shown in Fig. 5. The energy shown in Fig. 5

Table 2: Accuracy and FPGA throughput for CIFAR-10. In the “Model” column, “Full”, “L-2”, “L-1”, “FP”, “FL” indicate full-precision DNN, LightNN-2, LightNN-1, Fixed-point DNN, and FLightNN, respectively. The subscript “ $xWyA$ ” indicates x bits for weights and y bits for activations. The FLightNN results are shown in bold face. We use subscript a and b to denote the two trained FLightNNs for each network. These notations also apply for Table 3, 4 and 5.

ID	Model	Accuracy (%)	Storage (MB)	Throughput (images/s)	Speedup
1	Full	86.36	0.31	3.2e2	1×
	L-2 _{8W8A}	86.17	0.08	2.2e3	7.0×
	L-1 _{4W8A}	84.82	0.04	4.5e3	14.4×
	FP _{4W8A}	85.09	0.04	3.3e3	10.5×
	FL_{1a}	85.70	0.04	4.8e3	15.0×
	FL_{1b}	85.91	0.06	4.0e3	12.6×
2	Full	91.70	2.8	1.4e2	1×
	L-2 _{8W8A}	91.64	0.7	1.6e3	11.5×
	L-1 _{4W8A}	91.15	0.4	2.7e3	19.0×
	FP _{4W8A}	91.17	0.4	1.5e3	10.7×
	FL_{2a}	91.36	0.4	2.8e3	18.9×
	FL_{2b}	91.48	0.7	1.9e3	13.0×
3	Full	92.85	18.5	1.3	1×
	L-2 _{8W8A}	92.72	4.6	10.2	7.8×
	L-1 _{4W8A}	91.93	2.3	39.2	30.2×
	FP _{4W8A}	92.23	2.3	19.8	15.2×
	FL_{3a}	92.59	2.3	39.2	30.2×
	FL_{3b}	92.62	3.3	27.2	21.0×

Table 3: Accuracy and FPGA throughput for SVHN.

ID	Model	Accuracy (%)	Storage (MB)	Throughput (images/s)	Speedup
4	Full	94.96	0.12	2.2e3	1×
	L-2 _{8W8A}	94.90	0.03	4.5e3	2.09×
	L-1 _{4W8A}	94.16	0.02	8.3e3	3.63×
	FP _{4W8A}	93.70	0.02	3.7e3	1.70×
	FL_{4a}	94.67	0.02	6.7e3	3.11×
	FL_{4b}	94.88	0.03	5.3e3	2.37×
5	Full	96.44	0.4	1.1e3	1×
	L-2 _{8W8A}	96.38	0.1	2.1e3	2.00×
	L-1 _{4W8A}	95.93	0.05	3.7e3	3.53×
	FP _{4W8A}	96.02	0.05	1.8e3	1.71×
	FL_{5a}	96.21	0.06	3.2e3	3.06×
	FL_{5b}	96.24	0.08	3.0e3	2.84×

only includes the computational energy consumption for the largest layer of each network. We can clearly observe that FLightNNs provide a more continuous Pareto front for LightNN-2 and LightNN-1, regardless of the network type (*i.e.*, VGG or ResNet), size and the datasets. Similar to the observation in Sec. 5.2, in some networks FLightNNs can achieve higher accuracy than LightNN-1 with lower computational energy cost.

6 DISCUSSION

Since FLightNNs customize the k_i for each filter, LightNN-1 and LightNN-2 can be considered as two special cases for FLightNNs.

Table 4: Accuracy and FPGA throughput for CIFAR-100.

ID	Model	Accuracy (%)	Storage (MB)	Throughput (images/s)	Speedup
6	Full	69.16	2.8	2.5e2	1×
	L-2 _{8W8A}	68.84	0.7	1.6e3	6.4×
	L-1 _{4W8A}	67.32	0.4	2.7e3	10.6×
	FP _{4W8A}	67.67	0.4	1.5e3	5.98×
	FL_{6a}	68.59	0.4	2.7e3	10.6×
	FL_{6b}	68.76	0.6	1.8e3	6.88×
7	Full	71.22	11.2	7.4e1	1×
	L-2 _{8W8A}	70.96	2.8	6.0e2	8.11×
	L-1 _{4W8A}	69.71	1.4	1.1e3	15.2×
	FP _{4W8A}	69.34	1.4	6.9e2	9.26×
	FL_{7a}	70.85	1.4	1.1e3	15.2×
	FL_{7b}	70.87	2.4	7.4e2	9.98×

Table 5: Top-5 Accuracy and FPGA throughput for ImageNet.

ID	Model	Accuracy (%)	Storage (MB)	Throughput (images/s)	Speedup
8	L-2 _{8W8A}	75.04	1.8	2.7e2	1×
	L-1 _{4W8A}	72.94	0.9	5.2e2	1.95×
	FL_{8a}	74.80	1.5	3.1e2	1.16×
	FL_{8b}	75.00	1.7	2.8e2	1.06×

Table 6: FPGA resource utilization for different quantized DNN models.

ID	Model	Maximum resource utilization				Speedup
		BRAM	DSP	FF	LUT	
7	Full	896	642	69,344	128,339	1×
	L-2 _{8W8A}	1,024	4	66,491	90,949	8.11×
	L-1 _{4W8A}	1,024	4	66,491	90,949	15.2×
	FP _{4W8A}	1,024	514	7,110	90,949	9.26×
	FL_{7a}	1,024	4	84,192	90,949	9.98×
	FL_{7b}	1,024	4	63,648	80,940	15.2×
8	L-2 _{8W8A}	832	16	3,156	38,022	1×
	L-1 _{8W8A}	800	16	3,084	36,906	1.95×
	FL_{8a}	800	16	5,070	36,098	1.16×
	FL_{8b}	800	16	6,272	37,406	1.06×
Available		1,090	900	437,200	218,600	

Therefore, the Pareto front created by the searched FLightNN solutions should be the upper bound for the front of LightNN-1 and LightNN-2 with varied parameter numbers. We test this hypothesis on CIFAR-100 dataset using networks with varied number of convolutional filters. As shown in Fig. 6, the accuracy-storage Pareto-front created by FLightNNs is consistently higher than the LightNNs. This indicates that instead of only filling in the Pareto front of LightNNs, FLightNNs can *push forward* the Pareto front, due to their larger design space. The proposed differentiable training algorithm optimizes both k_i and weight values in an end-to-end fashion, and therefore significantly reduces searching effort compared to exhaustive or heuristic methods with multiple rounds of

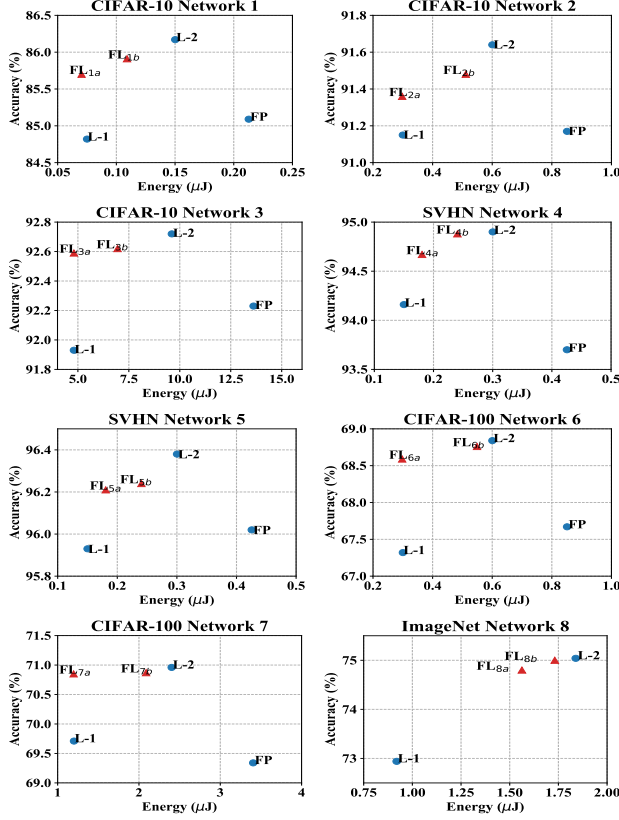


Figure 5: Accuracy and computational energy consumption in ASIC for different quantized models on CIFAR-10, SVHN, CIFAR-100 and ImageNet datasets. FLightNNs are marked as red triangles, while the other models are shown as blue dots.

training. Future work will further improve training efficiency by using optimized training loss [27] or proper labels [28].

7 CONCLUSION

In this paper, we propose FLightNNs which customize the number of shift operations for each filter of LightNNs. Equipped with the proposed differentiable training algorithm, FLightNNs can achieve a flexible trade-off between accuracy and speed/energy. Our experimental results on FPGA and ASIC simulations show that FLightNNs can provide a more continuous Pareto front for LightNN models and consistently outperform fixed-point DNNs *w.r.t.* both accuracy and speed/energy. Moreover, due to the gradual quantization nature of the differentiable training, FLightNNs can achieve higher accuracy than LightNN-1 without sacrificing speed and energy efficiency, and thus, push forward the Pareto-optimal front. These promising results suggest the potentials for FLightNNs to achieve fast and accurate inference on learning-based customized hardware.

ACKNOWLEDGMENTS

This research was supported in part by NSF CCF Grant No. 1815899.

REFERENCES

- [1] “Nnapi,” <https://developer.android.com/ndk/guides/neuralnetworks/>, accessed: 2018-10-15.
- [2] T.-J. Yang, Y.-H. Chen, and V. Sze, “Designing energy-efficient convolutional neural networks using energy-aware pruning,” in *IEEE CVPR*, pp. 6071–6079, 2017.
- [3] D. Lin, S. Talathi, S. Talathi, and S. Annapureddy, “Fixed point quantization of deep convolutional networks,” in *International Conference on Machine Learning*, pp. 2849–2858, 2016.

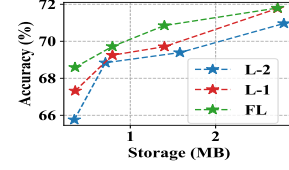


Figure 6: Accuracy-storage front for LightNN-2, LightNN-1 and FLightNN. The Pareto front of FLightNN is the upper bound of LightNNs.

- [4] R. Ding, Z. Liu, R. Shi, D. Marculescu, and R. Blanton, “Lightnn: Filling the gap between conventional deep neural networks and binarized networks,” in *Proceedings of the on Great Lakes Symposium on VLSI*, pp. 35–40, 2017.
- [5] R. Ding, Z. Liu, R. Blanton, and D. Marculescu, “Lightening the load with highly accurate storage- and energy-efficient lightnns,” *ACM transactions on Reconfigurable Technology and Systems*, vol. 11, no. 3, pp. 20–44, 2018.
- [6] I. Hubara, M. Courbariaux, D. Soudry, R. El-Yaniv, and Y. Bengio, “Binarized neural networks,” in *Advances in Neural Information Processing Systems*, pp. 4107–4115, 2016.
- [7] Y.-H. Chen, T. Krishna, J. S. Emer, and V. Sze, “Eyeriss: An energy-efficient reconfigurable accelerator for deep convolutional neural networks,” *IEEE Journal of Solid-State Circuits*, vol. 52, no. 1, pp. 127–138, 2017.
- [8] H. Zhao and J. Liu, “Processing-in-memory for energy-efficient neural network training: A heterogeneous approach,” in *IEEE/ACM International Symposium on Microarchitecture*, 2018.
- [9] S. Han, J. Pool, J. Tran, and W. Dally, “Learning both weights and connections for efficient neural network,” in *Advances in Neural Information Processing Systems*, pp. 1135–1143, 2015.
- [10] W. Wen, C. Wu, Y. Wang, Y. Chen, and H. Li, “Learning structured sparsity in deep neural networks,” in *Advances in Neural Information Processing Systems*, pp. 2074–2082, 2016.
- [11] D. M. Brooks, “Co-designed systems for deep learning hardware accelerators,” in *IEEE VLSI Design, Automation and Test (VLSI-DAT)*, *International Symposium*, pp. 1–1, 2018.
- [12] M. Sandler, A. Howard, M. Zhu, A. Zhmoginov, and L.-C. Chen, “Mobilenetv2: Inverted residuals and linear bottlenecks,” in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pp. 4510–4520, 2018.
- [13] H. Tann, S. Hashemi, R. I. Bahar, and S. Reda, “Hardware-software codesign of accurate, multiplier-free deep neural networks,” in *54th Design Automation Conference (DAC)*, pp. 1–6, 2017.
- [14] S. Gupta, A. Agrawal, K. Gopalakrishnan, and P. Narayanan, “Deep learning with limited numerical precision,” in *Proceedings of the 32nd International Conference on Machine Learning*, pp. 1737–1746, 2015.
- [15] S. Zhou, Y. Wu, Z. Ni, X. Zhou, H. Wen, and Y. Zou, “Dorefa-net: Training low bitwidth convolutional neural networks with low bitwidth gradients,” *arXiv preprint arXiv:1606.06160*, 2016.
- [16] M. Courbariaux, Y. Bengio, and J.-P. David, “Binaryconnect: Training deep neural networks with binary weights during propagations,” in *Advances in neural information processing systems*, pp. 3123–3131, 2015.
- [17] H. Liu, K. Simonyan, and Y. Yang, “Darts: Differentiable architecture search,” *arXiv preprint arXiv:1806.09055*, 2018.
- [18] C. Louizos, M. Welling, and D. P. Kingma, “Learning sparse neural networks through l₀ regularization,” in *International Conference on Learning Representations*, 2018.
- [19] Y. Bengio, N. Léonard, and A. Courville, “Estimating or propagating gradients through stochastic neurons for conditional computation,” *arXiv preprint arXiv:1308.3432*, 2013.
- [20] J. Han and C. Moraga, “The influence of the sigmoid function parameters on the speed of backpropagation learning,” in *International Workshop on Artificial Neural Networks*, pp. 195–201, 1995.
- [21] A. L. Maas, A. Y. Hannun, and A. Y. Ng, “Rectifier nonlinearities improve neural network acoustic models,” in *Proc. ICML*, vol. 30, no. 1, pp. 3, 2013.
- [22] D. P. Kingma and J. Ba, “Adam: A method for stochastic optimization,” *ICLR*, 2015.
- [23] C. Zhang, P. Li, G. Sun, Y. Guan, B. Xiao, and J. Cong, “Optimizing fpga-based accelerator design for deep convolutional neural networks,” in *Proceedings of the International Symposium on Field-Programmable Gate Arrays*, pp. 161–170, 2015.
- [24] Xilinx, “Vivado high-level synthesis,” <https://www.xilinx.com/products/design-tools/vivado/integration/esl-design.html>, 2017.
- [25] Y. Dong, R. Ni, J. Li, Y. Chen, J. Zhu, and H. Su, “Learning accurate low-bit deep neural networks with stochastic quantization,” *BMVC*, 2017.
- [26] S. M. U. Manual, “Synopsis inc,” *Mountain View, CA*, 2010.
- [27] R. Ding, T.-W. Chin, D. Marculescu, and Z. Liu, “Regularizing activation distribution for training binarized deep networks,” in *IEEE CVPR*, IEEE, 2019.
- [28] Z. Chen, R. Ding, T.-W. Chin, and D. Marculescu, “Understanding the impact of label granularity on cnn-based image classification,” in *2018 IEEE International Conference on Data Mining Workshops (ICDMW)*, IEEE, 2018, pp. 895–904.