

Regularizing Activation Distribution for Training Binarized Deep Networks

Ruizhou Ding, Ting-Wu Chin, Zeye Liu, Diana Marculescu
 Carnegie Mellon University

{rding, tingwuc, zeyel, dianam}@andrew.cmu.edu

Abstract

Binarized Neural Networks (BNNs) can significantly reduce the inference latency and energy consumption in resource-constrained devices due to their pure-logical computation and fewer memory accesses. However, training BNNs is difficult since the activation flow encounters degeneration, saturation, and gradient mismatch problems. Prior work alleviates these issues by increasing activation bits and adding floating-point scaling factors, thereby sacrificing BNN's energy efficiency. In this paper, we propose to use distribution loss to explicitly regularize the activation flow, and develop a framework to systematically formulate the loss. Our experiments show that the distribution loss can consistently improve the accuracy of BNNs without losing their energy benefits. Moreover, equipped with the proposed regularization, BNN training is shown to be robust to the selection of hyper-parameters including optimizer and learning rate.

1. Introduction

Recent years have witnessed tremendous success of Deep Neural Networks (DNNs) in various applications of image, video, speech, natural language, *etc* [16, 31]. However, the increased computation workload and memory access count required by DNNs pose a burden on latency-sensitive applications and energy-limited devices. Since latency and energy consumption are highly related to *computation* cost and *memory* access count, there has been a lot of research on reducing these two important design metrics [11, 41, 18, 6]. Binarized Neural Networks (BNNs) [22] that constrain the network weights and activations to be ± 1 have been proven highly efficient on custom hardware [51]. We also show later in Sec. 3.1 that a typical block of a BNN can be implemented in hardware with merely a few logical operators including XNOR gates, counters and comparators, and therefore greatly reduce the energy consumption and circuit area, as shown in Table 1.

In addition to the computational benefit brought by making the whole network binarized, another benefit of BNNs

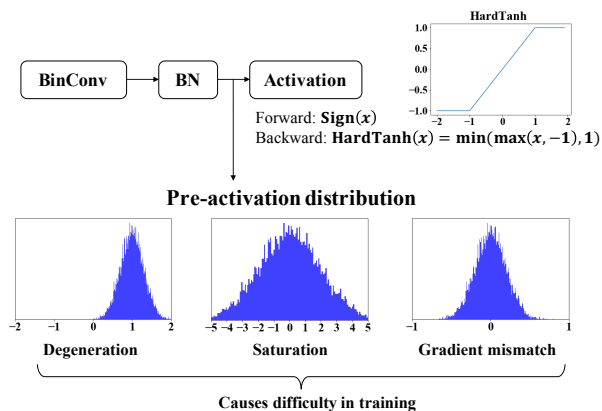


Figure 1: The basic Conv-BN-Act structure for BNN (BinConv: binary convolution; BN: batch normalization). The pre-activation distribution may exhibit from degeneration, saturation or gradient mismatch problem that causes difficulty in training.

is the huge reduction of memory footprint due to their 1-bit weights and activations. Prior work on extremely low-bit DNNs [9, 12, 28, 4, 46, 10] mainly focuses on few-bit weights and uses more bits for activations, while only a few [22, 29] target 1-bit weights *and* activations. However, reading and writing intermediate results (activations) generate a larger memory footprint than the weights [35]. For example, in the inference phase of a full-precision (32-bit) AlexNet with batch size 32, 92.7% of the memory footprint is caused by activations, while only 7.3% is caused by weights [35]. Therefore, the memory footprint of BNNs is significantly reduced due to their binary activations.

However, training accurate BNNs requires careful hyper-parameter selection [1], which makes the process more difficult than for their full-precision counterparts. Prior work has shown that this difficulty arises from the bounded activation function and the gradient approximation of the non-differentiable quantization function [4]. Even for full-precision DNNs, bounded activation functions (*e.g.*, *Sigmoid* or *Tanh*) usually lead to lower accuracy compared to the unbounded ones (*e.g.*, *ReLU*, *leaky ReLU*, or *SELU*)

due to the gradient vanishing problem [13, 5]. For binarized networks, a bounded activation (*i.e.*, *Sign* function) is used to lead to binary activations, and the *HardTanh* activation function is commonly used for gradient approximation [22, 38, 46]. As shown in Fig. 1, these bounded activation functions bring the following challenges (we use a convolutional layer as an example for illustration purposes): (i) *Degeneration*: If almost all the pre-activations of a channel have the same sign, then this channel will output nearly constant activations. In an extreme case, this channel degenerates to a constant. (ii) *Saturation*: If most of the pre-activations of a channel have a larger absolute value than the *HardTanh* threshold (*i.e.*, $|a| \geq 1$), then the gradients for these pre-activations will be zero. (iii) *Gradient mismatch*: If the absolute values of pre-activations are consistently smaller than the threshold (*i.e.*, $|a| < 1$), then this is equivalent to using a straight-through estimator (STE) for gradient computation [3]. While the STE generally performs well in computing gradients of staircase functions when training fixed-point DNNs, using STE for computing the gradient of *Sign* function causes larger approximation error than staircase function, and therefore causes worse gradient mismatch [4].

Due to the difficulty of BNN training, prior work along this track has traded the benefit of extremely-low energy consumption for higher accuracy. Hubara *et al.* largely increase the number of filters/neurons per convolutional/fully-connected layer [22]. Thus, while a portion of filters/neurons are blocked due to degeneration or gradient saturation, there is still a large *absolute* number of filters/neurons that can work well. Similarly, Mishra *et al.* also increase the width of the network to keep the BNN accuracy high [35].

In addition to increasing the number of network parameters, lots of work sacrifices BNNs’ pure-logical advantage by relaxing the precision constraint. Rastegari *et al.* approximate a full-precision convolution by using a binary convolution followed by a floating-point element-wise multiplication with a scaling matrix. [38]. Tang *et al.* use multiple-bit binarization for activations, which requires floating-point operators to compute the mean and residual error of activations [46]. Lin *et al.* approximate each filter and activation map using a weighted sum of multiple binary tensors [29]. All these approaches use scaling factors for weights and activations, making fixed-point multiplication and addition necessary for hardware implementation. Liu *et al.* added skip connections with floating point computations to the model [30]. While the models resulting from these approaches use XNOR convolution kernels, the extra multiplications and additions are not negligible. As shown in Table. 2, the energy cost for a typical convolutional layer of BNN is lower than the other binarized DNNs. The layer setting and the proposed approach for energy cost estima-

Table 1: Computational energy consumption and circuit area for different computation operators using a commercial 65nm process design kit [44]. The multiplier and adder are both 16-bit fixed-point operators.

	Energy (pJ)	Relative cost	Area (μm^2)	Relative cost
XNOR	7.6×10^{-4}	1×	4.2	1×
Counter	7.8×10^{-4}	10×	52	12×
Comparator	1.1×10^{-2}	14×	52	12×
Multiplier	1.6	2109×	3.0×10^3	718×
Adder	4.8×10^{-2}	64×	1.6×10^2	37×

Table 2: Computational energy for a convolutional layer with different types of binarizations.

	Pure-logical	Energy (μJ)	Relative cost
BNN [22]	Yes	1.42	1×
XNOR-Net [38]	No	4.34	3×
ABC-Net [29]	No	24.6	17×

tion are introduced in Appendix 6.1. Furthermore, since the hardware implementation of BNNs do not require digital signal processing (DSP) units, they greatly save circuit area and thus, can benefit IoT applications that have stringent area constraint [14, 8].

In this paper, we propose a general framework for activation regularization to tackle the difficulties encountered during BNN training. While prior work on weight initialization [13] and batch normalization [24] also regularizes activations, it does not address the challenges mentioned earlier for BNNs, as detailed in Sec. 4.2. Instead of regularizing the activation distribution in an implicit fashion as done in prior work [13, 24], we shape the distribution explicitly by embedding the regularization in the loss function. This regularization is shown to effectively alleviate the challenges for BNNs, and consistently increase the accuracy. Specifically, adding the distribution loss can improve the Top-1 accuracy of BNN AlexNet [22] on ImageNet from **36.1%** to **41.3%**, and improve the binarized wide AlexNet [35] from **48.3%** to **53.8%**. In summary, this paper has the following key contributions:

(i) To the best of our knowledge, we are the first to propose a framework for explicit activation regularization for binarized networks that consistently improve the accuracy.

(ii) Empirical results show that the proposed distribution loss is robust to the selection of training hyper-parameters. Code is available at: <https://github.com/ruizhoud/DistributionLoss>.

2. Related Work

Prior work has proposed various approaches to regularize the activation flow of full-precision DNNs, mainly to ad-

dress the gradient vanishing or exploding problem. Ioffe *et al.* propose batch normalization to centralize the activation distribution, accelerate training, and achieve higher accuracy [24]. Similarly, Huang *et al.* normalize the weights with zero mean and unit norm followed by scaling factors [20]. Shang *et al.* extend the normalization idea to residual networks using normalized propagation [42], while Ba *et al.* and Salimans *et al.* normalize the activations of Recurrent Neural Network (RNN) by layer-wise normalization and weight reparameterization, respectively [2, 40]. In addition, some prior work develops good initialization strategy to regularize the activations in the initial state [33, 48], or proposes new activation functions to maintain stable activation distribution across layers [27, 32].

However, these approaches on full-precision networks do not address the difficulty of training networks with binarized activations. Prior work on binarized DNNs alleviates this problem mainly by approximating the full-precision activations with multiple-bit representations and floating-point scaling factors [46, 4, 35, 37, 34, 12, 29, 19]. Tang *et al.* introduced scaling layers and use 2 bits for activations [46]. Cai *et al.* use multi-level activation function for inference and variants of ReLU for gradient computation to reduce gradient mismatch [4]. Polino *et al.* leverage knowledge distillation to guide training and improve the accuracy with multiple bits for activations [37]. Lin *et al.* approximate both weights and activations with multiple binary bases associated with floating-point coefficients [29]. While these approaches improve the accuracy for binarized networks, they sacrifice the energy efficiency due to the increased bits and the required DSP units for the additions and multiplications.

3. Activation Regularization

In this section, we first show that BNN blocks can be implemented with pure-logical operators in hardware while the other binarized networks (including XNOR-Net and ABC-Net) based on scaling factors require additional full-precision operations. Then, we propose a framework to address the problems of activation and gradient flow incurred in the training process of BNNs. Finally, we discuss the effectiveness of this framework.

3.1. Binarized DNNs

Binarized DNNs constrain the weights and activations to be ± 1 , making the convolution between weights and activations use only *xnor* and *count* operators. In this subsection, we introduce the structure of three typical binarized DNNs, and analyze their hardware implication.

BNN: As shown in Fig 2, the basic block for BNN [22] is composed of a binary convolution, a batch normalization (BN) and an optional max pooling layer, followed by a sign activation function. Without changing the input-output

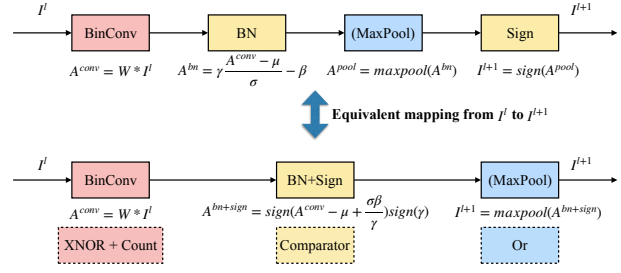


Figure 2: Basic block for convolutional BNN [22]. The activations I^l and weights W are binarized to ± 1 . The inference of this block can be implemented on hardware with only logical operators.

mapping of this block, we can reorder the max pooling layer and the sign function, and then, combine the BN layer and sign function to be a comparator of the convolution results A^{conv} and input-independent variables $\mu + \frac{\sigma\beta}{\gamma}$, where μ and σ are the moving mean and variance of per-channel activations, which are obtained from training data and fixed in the testing phase; β and γ are trainable parameters in the BN layer. Therefore, the inference of this BNN block can be implemented in hardware with pure-logical operators. This transformation can also be applied to binarized fully-connected layers followed by BN, pooling and sign function.

XNOR-Net: Different from BNN, XNOR-Net [38] approximates the activations A^{bn} after the BN layer with their signs and scaling factors computed by the average of the absolute values of these activations, as shown in Fig. 9 in Appendix. Since the scaling factors are input-dependent, the full-precision multiplications and additions cannot be eliminated.

ABC-Net: ABC-Net [29], shown in Fig. 10 in Appendix, approximates both weights and activations with a linear combination of pre-defined bases, and therefore making the convolution kernel binarized. However, the approximation prior to binary convolution and the scaling operations after the convolution require full-precision multiplications and additions that cannot be eliminated.

3.2. Regularizing Activation Distribution

In this section we first introduce some notations and formally define the difficulties encountered when training BNNs. We denote $A^{b,l,c}$ as the pre-activations (activations prior to the *Sign* function) for the c -th channel of the l -th layer for the b -th batch of data. Thus, $A^{b,l,c}$ is a 3D tensor with size $B \times W \times H$ where B is the batch size, W and H are the width and height of the activation map. From this point on, to avoid clutter, we will omit the superscript of A whenever possible. $A_{(q)}$ denotes the q quantile of A 's elements where $0 \leq q \leq 1$. We define degeneration, satu-

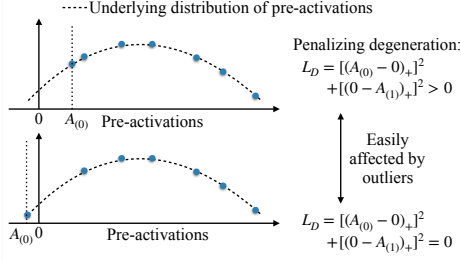


Figure 3: Motivation for adjusting regularization. The loss function directly formulated from hypothesis (e.g., degeneration) relies on the minimum (or maximum) of the pre-activations, and therefore is sensitive to outliers.

ration, and gradient mismatch as follows:

$$\begin{aligned} \textbf{Degeneration: } & A_{(0)} \geq 0 \text{ or } A_{(1)} \leq 0 \\ \textbf{Saturation: } & |A|_{(0)} \geq 1 \\ \textbf{Gradient mismatch: } & |A|_{(1)} \leq 1 \end{aligned} \quad (1)$$

where $|A|_{(q)}$ is the q quantile for $|A|$, and we use 1 because it is the threshold of *HardTanh* activation shown in Fig. 1.

To alleviate the aforementioned problems, we propose to add the distribution loss in the objective function to regularize the activation distribution. Using degeneration as an example, an intuitive way of formulating a loss to avoid the degeneration problem for A is $L_D = [(A_{(0)} - 0)_+]^2 + [(0 - A_{(1)})_+]^2$, where $(\cdot)_+$ is the ReLU function. However, this may lead to too loose regularization since a small outlier can make this loss zero, as shown in Fig. 3. In addition, this formulation of L_D is not differentiable *w.r.t.* the pre-activations A .

Therefore, we propose a three-stage framework consisting of hypothesis formulation, adjusting regularization, and enabling differentiability, to systematically formulate an outlier-robust and differentiable regularization, as shown in Fig. 4. First, based on the prior hypothesis about the activation distribution, we can formulate a loss function to penalize the unwanted distribution. Then, if this formulation uses large-variance estimators (e.g., maximum or minimum of samples), we can use relaxed estimators (e.g., quantiles) to increase robustness to outliers. Finally, if the formulated loss is not differentiable, we need to approximate it by assuming the type of parametric distribution (e.g., Gaussian), and approximate the non-differentiable estimators with the distribution parameters.

Degeneration. We first formulate the degeneration hypothesis in the loss function as $L_D = [(A_{(0)} - 0)_+]^2 + [(0 - A_{(1)})_+]^2$. To make the loss function more robust to outliers, we adjust the regularization by using relaxed quantiles ϵ and $1 - \epsilon$, with $L_D = [(A_{(\epsilon)} - 0)_+]^2 + [(0 - A_{(1-\epsilon)})_+]^2$. Then, to make L_D differentiable so that it can fit in the backpropagation training, we first assume a parameterized distribution

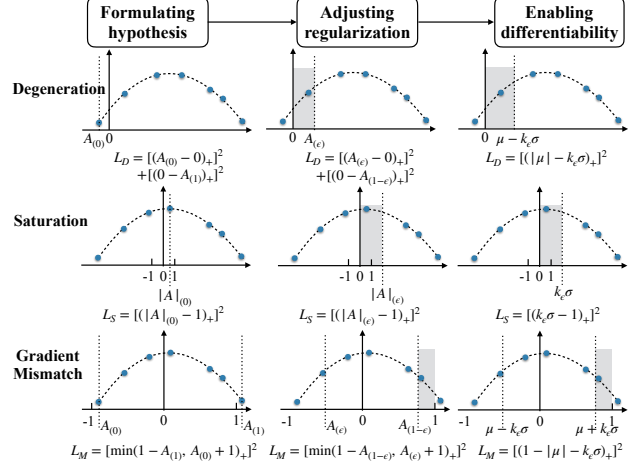


Figure 4: Proposed framework for formulating the differentiable loss function to regularize activation distribution. Starting from the three hypotheses (“degeneration”, “saturation” and “gradient mismatch”), we can formulate the loss function L_D , L_S and L_M for them, respectively. We omit the superscript for A and L for better representation.

for the pre-activations A and then use its parameters to formulate a differentiable L_D . Based on the heuristics from prior art [28, 4, 36], we assume that the values of A follow a Gaussian distribution $\mathcal{N}(\mu, (\sigma)^2)$, where the μ and σ can be estimated by the sample mean and standard deviation over the 3D tensor. Thus, we can formulate the ϵ quantile by $\mu - k_\epsilon \sigma$ where k_ϵ is a constant determined by ϵ . Therefore, $L_D = [(\mu - k_\epsilon \sigma - 0)_+]^2 + [(0 - (\mu + k_\epsilon \sigma))_+]^2 = [(|\mu| - k_\epsilon \sigma)_+]^2$.

Saturation. The saturation problem can be penalized by $L_S = [(|A|_{(0)} - 1)_+]^2$, where $|A|_{(0)}$ is the minimum value of $|A|$. By adjusting the regularization, we have $L_S = [(|A|_{(\epsilon)} - 1)_+]^2$. Since L_D already eliminates the degeneration problem, we find that simply assuming A has a zero mean (i.e., $\mathcal{N}(0, (\sigma)^2)$) works well empirically. Thus, the loss function is formulated as $L_S = [(k_\epsilon \sigma - 1)_+]^2$.

Gradient mismatch. When most of the activations lie in the range of $[-1, 1]$, the backward pass is simply using a STE for the gradient computation of the sign function, causing the gradient mismatch problem. Therefore, we can formulate the loss as $L_M = [\min(1 - A_{(1)}, A_{(0)} + 1)_+]^2$. Similarly, relaxing the regularization leads to $L_M = [\min(1 - A_{(1-\epsilon)}, A_{(\epsilon)} + 1)_+]^2$. With a Gaussian assumption, we have $L_M = [\min(1 - \mu - k_\epsilon \sigma, \mu - k_\epsilon \sigma + 1)_+]^2 = [(1 - |\mu| - k_\epsilon \sigma)_+]^2$.

Then, in the training phase, we add the distribution loss for the b -th batch of input data:

$$L_{DL}^b = \sum_{l,c} L_{DL}^{b,l,c} = \sum_{l,c} L_D^{b,l,c} + L_S^{b,l,c} + L_M^{b,l,c}, \quad (2)$$

and the total loss for b -th batch is:

$$L_{total}^b = L_{CE}^b + \lambda L_{DL}^b \quad (3)$$

where L_{CE}^b is the cross-entropy loss, and λ is a coefficient to balance the losses.

3.3. Intuition for the Proposed Distribution Loss

The distribution loss is proposed to alleviate training problems for pure-logical binarized networks. In contrast with full-precision networks, BNNs use a bounded activation function and therefore exhibit the gradient saturation and mismatch problems. By regularizing the activations, the distribution loss maintains the effectiveness of the back-propagation algorithm, and thus, can speedup training and improve the accuracy.

Since the distribution loss changes the optimization objective, one concern may be that it will lead to a configuration far from the global optimal of the cross-entropy loss function. However, prior theoretical [43, 25, 7] and empirical [23] work has shown that a deep neural network can have many high-quality local optima. Kawaguchi proved that under certain conditions, every local minimum is a global minimum [25]. Through experiments, Im *et al.* show that using different optimizers, the achieved local optima are very different [23]. These insights show that adding the distribution loss may deviate the training away from the original optimal, but can still lead to a new optimal with high accuracy. Moreover, the distribution loss diagnoses the poor conditions of the activation flow, and therefore may achieve higher accuracy. Our experiment results confirm this hypothesis.

4. Experimental Results

In this section, we first evaluate the accuracy improvement by the proposed distribution loss on CIFAR-10, SVHN, CIFAR-100 and ImageNet. Then, we visualize the histograms of the regularized activation distribution. Finally, we analyze the robustness of our approach to the hyper-parameter selection.

4.1. Accuracy Improvement

Training configuration. We use fully convolutional VGG-style networks for CIFAR-10 and SVHN, and ResNet for CIFAR-100. All of them use the ADAM optimizer [26] as suggested by Hubara *et al.* [22]. For the BNN trained with distribution loss (BNN-DL), we compute the loss with the activations prior to each binarized activation function (*i.e.*, *Sign* function that uses *HardTanh* for gradient computation). Unless noted otherwise, we set the coefficient k_ϵ to be 1, 0.25 and 0.25 for L_D , L_S and L_M , respectively, and set λ to be 2. To show the statistical significance, all the experiments for CIFAR-10 and SVHN are averaged over five

experiments with different parameter initialization seeds. The details of the network structure and training scheme for each dataset is as follows:

CIFAR-10. The network structure can be formulated as: $x\text{C}-x\text{C}-\text{MP}-2x\text{C}-2x\text{C}-\text{MP}-4x\text{C}-4x\text{C}-10\text{C}-\text{GP}$, where $x\text{C}$ indicates a convolutional layer with x filters, MP and GP indicate max pooling and global pooling layers, respectively. 3×3 filter size is used for all the convolutional layers. We vary the x to different values ($\{128, 179, 256, 384\}$) to explore the trade-off between accuracy and energy cost, which are shown in Table 3 as networks 2-5. We also train a small BNN without the two $4x\text{C}$ layers for CIFAR-10, which is network 1 in Table 3. Each convolutional layer has binarized weights and is followed by a batch normalization layer and a sign activation function. The learning rate schedule follows the code from BNN authors [21]. All networks are trained for 200 epochs.

SVHN. The network structure for SVHN is the same as CIFAR-10, except that the x is varied from $\{51, 64, 96, 128\}$, shown by networks 6-9 in Table 3. The initial learning rate value is $1e-2$, and decays by a factor of 10 at epochs 20, 40 and 45. We train 50 epochs in total.

CIFAR-100. We use the *full pre-activation* variant of ResNet [17] with 20 layers for CIFAR-100. Prior work has shown the difficulty of training ResNet-based BNNs without scaling layers [49]. Since in ResNet-based BNNs the main path activations and residual path activations do not have matching scales, directly adding the two activations will cause difficulty in training. Therefore, we add two batch normalization layers after these two activations to maintain stable activation scales. Similar to CIFAR-10 and SVHN, we also vary the number of filters per layer. Details on the network structure are included in Appendix 6.3. The initial learning rate is set to $1e-4$, and reduced by a factor of 3 every 80 epochs. The networks are trained for 300 epochs.

Results on CIFAR-10, SVHN and CIFAR-100. As shown in Table 3, the accuracy for BNN-DL is consistently higher than the baseline BNN. The accuracy gap between BNN and BNN-DL is generally larger than their standard deviations. Using t-test, the p-values for all the network 1-9 are smaller than 0.005, which demonstrates the statistical significance of our improvements. In addition to accuracy results, we also show the computational energy cost for each network, obtained by summing up the energy of each operation for the inference of a single input image. Note that this cost excludes the energy of memory accesses, which is the same for BNN and BNN-DL in the inference phase. We also visualize the trade-off between accuracy and energy cost in Fig. 5. In most cases, the BNN-DL with a smaller model size can achieve the same or higher accuracy than the BNN with a larger size.

The use of the distribution loss improves the testing accuracy mostly because it regularizes the activation and gra-

Table 3: Accuracy improvement with distribution loss. Network depth is defined as the number of convolutional layers, while the network width is defined as the number of filters in the largest layer. The best results are shown in bold face. *All the accuracy for CIFAR-10 and SVHN is averaged over five experiments with different weight initialization.*

Dataset	Network ID	Depth/width	Params storage	Energy cost (μ J)	Accuracy (mean $\pm$ std) (%)	
					BNN	BNN-DL
CIFAR-10	1	5/256	0.4 MB	0.30	80.61 $\pm$ 0.49	83.33 $\pm$ 0.32
	2	7/512	0.6 MB	0.47	87.54 $\pm$ 0.38	89.13 $\pm$ 0.23
	3	7/716	1.1 MB	0.93	88.99 $\pm$ 0.13	90.28 $\pm$ 0.28
	4	7/1024	2.3 MB	1.89	90.09 $\pm$ 0.10	91.01 $\pm$ 0.09
	5	7/1536	3.8 MB	4.23	90.68 $\pm$ 0.11	91.56 $\pm$ 0.16
SVHN	6	7/204	0.09 MB	0.08	96.23 $\pm$ 0.15	96.57 $\pm$ 0.12
	7	7/256	0.15 MB	0.12	96.53 $\pm$ 0.11	96.95 $\pm$ 0.10
	8	7/384	0.3 MB	0.27	97.15 $\pm$ 0.15	97.34 $\pm$ 0.05
	9	7/512	0.6 MB	0.47	97.34 $\pm$ 0.07	97.51 $\pm$ 0.03
CIFAR-100	10	20/1024	5.6 MB	53.7	60.40	68.17
	11	20/1536	12.6 MB	120.9	64.57	71.53
	12	20/2048	22.3 MB	215.0	66.07	73.42

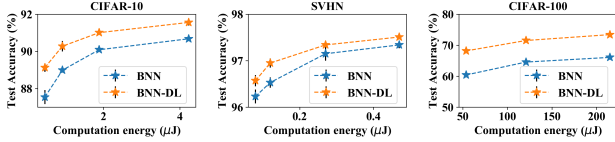


Figure 5: Accuracy and energy Pareto-optimal curve for CIFAR-10, SVHN and CIFAR-100. The error bars for CIFAR-10 and SVHN show the standard deviation of testing accuracy.

dient flow in the training phase, so that the networks can better fit the dataset. As shown in Fig. 6, the training loss for BNN-DL is consistently lower than the BNN baseline after a few epochs. For most of the experiments, distribution loss is found to converge to a very small number (*e.g.*, 1/10000 of the initial value) in the first few epochs. This indicates that the network can be easily regularized by the distribution loss, which then improves the rest of the training process.

Comparison with prior art. We also compare our results with prior work on binarized networks as shown in Table 4. For CIFAR-10 and SVHN, we follow the same network configuration by Hou *et al.*, and also split the dataset into training, validation and testing sets as they do [19]. Table 4 shows that by just applying the distribution loss when training BNNs can achieve higher accuracy than the baseline BNN [22], XNOR-Net [38] and LAB [19]. We also show the normalized energy cost for the models. Since XNOR-Net and LAB use scaling factors for the weights and activations, which introduces the need for full-precision operations, XNOR-Net and LAB require $4.5\times$ energy cost than BNN and BNN-DL. We use 16-bit fixed-point multipliers and adders instead of 32-bit floating-point opera-

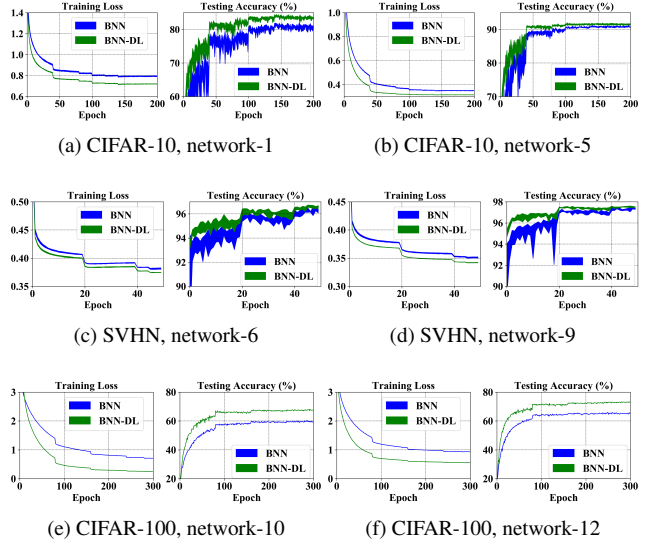


Figure 6: Training loss and testing accuracy curves for different networks with or without distribution loss. The widths of the curves for CIFAR-10 and SVHN are 2 standard deviation ranges.

tors to estimate the energy cost of these full-precision operations because prior quantization work shows that 16-bit fixed-point operation is generally enough for maintaining accuracy [15]. For CIFAR-100, the closest work that reports 1-bit weights and low-bit activations is by Polino *et al.* [37], where they use a 7.9MB ResNet with 2-bit activations for CIFAR-100, which presumably has larger energy cost than our 5.6MB model with 1-bit activations, and our results on accuracy surpass theirs by a large margin.

Results on ImageNet. Having shown the effectiveness of the distribution loss on small datasets, we extend our analy-

Table 4: Comparison with prior art using 1-bit weights and activations, in terms of accuracy and computation energy on different datasets. The best results are shown in bold face.

Dataset	Model	Pure-logical	Energy cost	Accuracy
CIFAR-10	BNN [22]	Yes	$1\times$	87.13%
	XNOR-Net [38]	No	$4.5\times$	87.38%
	LAB [19]	No	$4.5\times$	87.72%
	BNN-DL	Yes	$1\times$	89.90%
SVHN	BNN [22]	Yes	$1\times$	96.50%
	XNOR-Net [38]	No	$4.5\times$	96.57%
	LAB [19]	No	$4.5\times$	96.64%
	BNN-DL	Yes	$1\times$	97.23%
CIFAR-100	BNN [22]	No	$1\times$	60.40%
	DQ-2bit [37]	No	-	49.32%
	BNN-DL	No	$1\times$	68.17%

Table 5: Robustness to the selection of optimizer, learning rate, and network structure. CIFAR-10 is used for illustrating the results.

	BNN	BNN-DL
Momentum	66.02%	89.37%
Nesterov	68.66%	89.22%
Adam	88.12%	89.62%
RMSprop	87.39%	90.24%
$lr_{init}=1e-1$	82.19%	89.60%
$lr_{init}=5e-3$	88.12%	89.62%
$lr_{init}=2e-4$	85.62%	88.73%
VGG	88.12%	89.62%
ResNet-18	85.71%	90.47%

sis to a larger image dataset - ImageNet ILSVRC-2012 [39]. We consider AlexNet, which is the most commonly adopted network in prior art on binarized DNNs [22, 38, 52, 46, 35]. We compare our BNN-DL with the baseline BNN [22], XNOR-Net [38], DoReFa-Net [52], Compact Net [46], and WRPN [35]. The BNN uses binarized weights for the whole network [22], while XNOR-Net and DoReFa-Net keep the first convolutional layer and last fully-connected layer with full-precision weights [38, 52]. Compact Net uses full-precision weights for the first layer but binarizes the last layer, and uses 2 bits for the activations [46]. WRPN doubles the filter number of XNOR-Net, and uses full-precision weights for both the first and last layers [35]. Also, BNN uses 64 and 192 filters while the other networks use 96 and 256 filters (or doubling these numbers as WRPN does) for the first two convolutional layers. We train our BNN-DL using the same settings as prior work, except that we use 1-bit activations instead of 2-bit when comparing with Compact Net. The learning rate policy follows prior implementations [21], but starts from 0.01. As shown in Table 6, BNN-DL consistently outperforms the accuracy of the baseline models. All baseline models except BNN use scaling factors to approximate activations while we keep them binarized. Therefore, our model also has lower energy cost than the prior models. In addition, we highlight that our BNN-DL can outperform Compact Net though we use fewer bits for activations.

4.2. Regularized Activation Distribution

To show the regularization effect of the distribution loss, we plot the distribution of the pre-activations for the baseline BNN and for our proposed BNN-DL. More specifically, we conduct inference for network 2 on CIFAR-10, and extract the (floating-point) activations right after the batch normalization layer prior to the binarized activation function of the fourth convolutional layer with 256 filters. There-

Table 6: Comparison with prior art on ImageNet with AlexNet-based topology. We use the same model structure as prior work, except that Compact Net uses 2 bits for activations while we only use 1 bit. Training with distribution loss outperforms prior work consistently.

Model	Baseline		Ours	
	Top-1	Top-5	Top-1	Top-5
BNN [22]	36.1%	60.1%	41.3%	65.8%
XNOR-Net [38]	44.2%	69.2%	47.8%	71.5%
DoReFa-Net [52]	43.5%	-	47.8%	71.5%
Compact Net [46]	46.6%	71.1%	47.6%	71.9%
WRPN [35]	48.3%	-	53.8%	77.0%

fore, for each of the 256 output channels, we get its values across the whole dataset. Then, for illustration purposes, we select four channels from baseline BNN and our proposed BNN-DL, respectively, and plot the histogram of these per-channel values, as shown in Fig. 7. The four channels' activation distributions for the baseline BNN are picked to show the degeneration, gradient mismatch, and saturation problems, while the distributions for BNN-DL are randomly selected. From Fig. 7a we can see that the good weight initialization strategy [13] and batch normalization [24] adopted for BNNs do not solve the distribution problems.

To show that BNN-DL alleviates these challenges, we compute the standard deviation of activations for each of these 256 channels, as well as their positive ratio, which is the proportion of positive values. As shown in Fig. 8, the standard deviation of BNN-DL is more regularized and centralized than that of BNN. The channel with very small standard deviation like the middle two histograms in Fig. 7a is rarely seen in BNN-DL, while BNN has a long tail in the area of small standard deviations. This indicates that without explicit regularization, the scale factors of batch normalization layer could shrink to very small values, causing the gradient mismatch problem. From Fig. 8, we can also

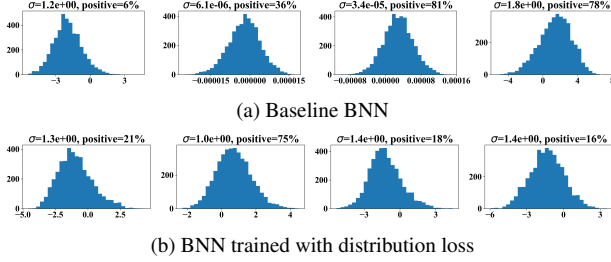


Figure 7: Activation distribution for BNN trained (a) without or (b) with distribution loss. Each histogram refers to the activations of one channel. In (a), the channel in the left histogram shows a generation problem, the middle two show gradient mismatch, and the right one shows saturation problem. σ is standard deviation, and “positive” refers to the ratio of positive activations.

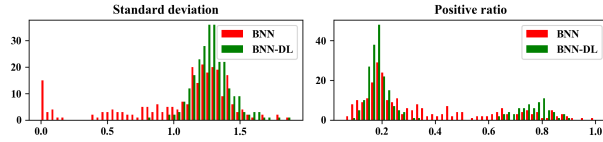


Figure 8: Histogram of standard deviation and positive ratio of per-channel activations.

observe that the positive ratio of BNN has more extreme values (*i.e.*, those close to 0 or 1) than BNN-DL. This indicates that the degeneration problem is reduced by distribution loss. Interestingly, we can see that the positive ratio of BNN-DL also deviates away from 0.5. We conjecture that this is because the activations centered at 0 are more prone to gradient mismatch, and thus, be penalized by L_M .

4.3. Robustness to Hyper-parameter Selection

Another benefit of the distribution loss is its robustness to the selection of the training hyper-parameters. Prior work [1] has shown that the accuracy of BNNs is sensitive to the training optimizer. We observe the same phenomenon by training BNNs with different optimizers including SGD with momentum, SGD with Nesterov [45], Adam [26] and RMSprop [47]. However, when training BNN with distribution loss, these optimizers can be consistently improved, as shown in Table 5. We use CIFAR-10 for the experiments in this subsection. We use the same weight decay and learning rate schedule as Zagoruyko *et al.* [50] for Momentum and Nesterov, and change the initial learning rate to $1e-4$ for RMSprop. We use the same setting as Hubara *et al.* [21] for Adam. Each model is trained for 200 epochs, and the best testing accuracy is reported. Then, we vary the learning rate schedule of Hubara *et al.*’s implementation [21] by scaling the learning rate at each epoch by a constant. Table 5 shows that BNN-DL is more robust to the selection of learning rate values. Furthermore, we show that the BNN-DL can work well for both VGG-style networks with stacked convolutional layers and ResNet-18 which in-

Table 7: Accuracy for BNN-DL on CIFAR-10 with varied regularization levels. $\lambda = 0$ indicates the baseline BNN.

λ	0	0.2	2	20	200	2000
Acc. (%)	87.39	90.12	90.16	90.18	90.61	90.20

cludes skip connections. The VGG-style network uses the network 2 in Table 3. The ResNet-18 structure uses the pre-activation variant [17] with added batch normalization layers as described in Sec. 4.1. Since BNN has non-regularized activations, maintaining the activation flow in the training process requires more careful picking of hyper-parameter values. However, the distribution loss applies regularization to the activations, making the network easier to train, and therefore reduces the sensitivity to hyper-parameter selection.

We also show that the distribution loss is robust to the selection of the introduced hyper-parameter, λ coefficient, which indicates the regularization level of the distribution loss. As shown in Table 7, by varying λ from 0.2 to 2000, the accuracy for BNN-DL is consistently higher than the baseline BNN. As mentioned in Sec. 4.1, the distribution loss quickly decays to a small magnitude in the first few epochs, and we find that this holds for a wide range of λ . The robustness analysis indicates that the distribution loss is a handy tool to regularize activations, without the need of much hyper-parameter tuning.

5. Conclusion

In this paper, we tackle the difficulty of training BNNs with 1-bit weights and 1-bit activations. The difficulty arises from the unregularized activation flow that may cause degeneration, saturation and gradient mismatch problems. We propose a framework to embed this insight into the loss function by formulating our hypothesis, adjusting regularization and enabling differentiability, and thus, explicitly penalizing the activation distributions that may lead to the training problems. Our experiments show that BNNs trained with the proposed distribution loss have regularized activation distribution, and consistently outperform the baseline BNNs. The proposed approach can significantly improve the accuracy of the state-of-the-art networks using 1-bit weights and activations for AlexNet on ImageNet dataset. In addition, this approach is robust to the selection of training hyper-parameters including learning rate and optimizer. These results show that distribution loss can generally benefit the training of binarized networks which enable latency and energy efficient inference on mobile devices.

Acknowledgement

This research was supported in part by NSF CCF Grant No. 1815899, and NSF award number ACI-1445606 at the Pittsburgh Supercomputing Center (PSC).

References

- [1] M. Alizadeh, J. Fernandez-Marques, N. D. Lane, and Y. Gal. A systematic study of binary neural networks' optimisation. In *International Conference on Learning Representations*, 2019.
- [2] J. L. Ba, J. R. Kiros, and G. E. Hinton. Layer normalization. *arXiv preprint arXiv:1607.06450*, 2016.
- [3] Y. Bengio, N. Léonard, and A. Courville. Estimating or propagating gradients through stochastic neurons for conditional computation. *arXiv preprint arXiv:1308.3432*, 2013.
- [4] Z. Cai, X. He, J. Sun, and N. Vasconcelos. Deep learning with low precision by half-wave gaussian quantization. *arXiv preprint arXiv:1702.00953*, 2017.
- [5] 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)*, pages 895–904. IEEE, 2018.
- [6] T.-W. Chin, R. Ding, and D. Marculescu. Adascale: Towards real-time video object detection using adaptive scaling. *arXiv preprint arXiv:1902.02910*, 2019.
- [7] A. Choromanska, M. Henaff, M. Mathieu, G. B. Arous, and Y. LeCun. The loss surfaces of multilayer networks. In *Artificial Intelligence and Statistics*, pages 192–204, 2015.
- [8] F. Conti, R. Schilling, P. D. Schiavone, A. Pullini, D. Rossi, F. K. Gürkaynak, M. Muehlberghuber, M. Gautschi, I. Loi, G. Haugou, et al. An iot endpoint system-on-chip for secure and energy-efficient near-sensor analytics. *IEEE Transactions on Circuits and Systems I: Regular Papers*, 64(9):2481–2494, 2017.
- [9] 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*, pages 3123–3131, 2015.
- [10] 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 (TRETS)*, 11(3):17, 2018.
- [11] 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 2017*, pages 35–40. ACM, 2017.
- [12] J. Faraone, N. Fraser, M. Blott, and P. H. Leong. Syq: Learning symmetric quantization for efficient deep neural networks. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 4300–4309, 2018.
- [13] X. Glorot and Y. Bengio. Understanding the difficulty of training deep feedforward neural networks. In *Proceedings of the thirteenth international conference on artificial intelligence and statistics*, pages 249–256, 2010.
- [14] G. Gobieski, N. Beckmann, and B. Lucia. Intelligence beyond the edge: Inference on intermittent embedded systems. *arXiv preprint arXiv:1810.07751*, 2018.
- [15] S. Gupta, A. Agrawal, K. Gopalakrishnan, and P. Narayanan. Deep learning with limited numerical precision. In *International Conference on Machine Learning*, pages 1737–1746, 2015.
- [16] K. He, X. Zhang, S. Ren, and J. Sun. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 770–778, 2016.
- [17] K. He, X. Zhang, S. Ren, and J. Sun. Identity mappings in deep residual networks. In *European conference on computer vision*, pages 630–645. Springer, 2016.
- [18] Y. He, J. Lin, Z. Liu, H. Wang, L.-J. Li, and S. Han. Amc: Automl for model compression and acceleration on mobile devices. In *Proceedings of the European Conference on Computer Vision (ECCV)*, pages 784–800, 2018.
- [19] L. Hou, Q. Yao, and J. T. Kwok. Loss-aware binarization of deep networks. In *International Conference on Learning Representations*, 2017.
- [20] L. Huang, X. Liu, Y. Liu, B. Lang, and D. Tao. Centered weight normalization in accelerating training of deep neural networks. In *Proceedings of the IEEE International Conference on Computer Vision*, pages 2803–2811, 2017.
- [21] I. Hubara. Binarynet.pytorch. https://github.com/itayhubara/BinaryNet.pytorch/blob/master/models/vgg_cifar10_binary.py, 2017.
- [22] I. Hubara, M. Courbariaux, D. Soudry, R. El-Yaniv, and Y. Bengio. Binarized neural networks. In *Advances in neural information processing systems*, pages 4107–4115, 2016.
- [23] D. J. Im, M. Tao, and K. Branson. An empirical analysis of deep network loss surfaces. 2016.
- [24] S. Ioffe and C. Szegedy. Batch normalization: Accelerating deep network training by reducing internal covariate shift. *arXiv preprint arXiv:1502.03167*, 2015.
- [25] K. Kawaguchi. Deep learning without poor local minima. In *Advances in Neural Information Processing Systems*, pages 586–594, 2016.
- [26] D. P. Kingma and J. Ba. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*, 2014.
- [27] G. Klambauer, T. Unterthiner, A. Mayr, and S. Hochreiter. Self-normalizing neural networks. In *Advances in Neural Information Processing Systems*, pages 971–980, 2017.
- [28] D. Lin, S. Talathi, and S. Annapureddy. Fixed point quantization of deep convolutional networks. In *International Conference on Machine Learning*, pages 2849–2858, 2016.
- [29] X. Lin, C. Zhao, and W. Pan. Towards accurate binary convolutional neural network. In *Advances in Neural Information Processing Systems*, pages 345–353, 2017.
- [30] Z. Liu, B. Wu, W. Luo, X. Yang, W. Liu, and K.-T. Cheng. Bi-real net: Enhancing the performance of 1-bit cnns with improved representational capability and advanced training algorithm. In *Proceedings of the European Conference on Computer Vision (ECCV)*, pages 722–737, 2018.
- [31] J. Lu, C. Xiong, D. Parikh, and R. Socher. Knowing when to look: Adaptive attention via a visual sentinel for image captioning. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, volume 6, page 2, 2017.
- [32] A. L. Maas, A. Y. Hannun, and A. Y. Ng. Rectifier nonlinearities improve neural network acoustic models. In *Proc. icml*, volume 30, page 3, 2013.

- [33] D. Mishkin and J. Matas. All you need is a good init. *arXiv preprint arXiv:1511.06422*, 2015.
- [34] A. Mishra and D. Marr. Apprentice: Using knowledge distillation techniques to improve low-precision network accuracy. In *International Conference on Learning Representations*, 2018.
- [35] A. Mishra, E. Nurvitadhi, J. J. Cook, and D. Marr. WRPN: Wide reduced-precision networks. In *International Conference on Learning Representations*, 2018.
- [36] E. Park, J. Ahn, and S. Yoo. Weighted-entropy-based quantization for deep neural networks. In *IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2017.
- [37] A. Polino, R. Pascanu, and D. Alistarh. Model compression via distillation and quantization. In *International Conference on Learning Representations*, 2018.
- [38] M. Rastegari, V. Ordonez, J. Redmon, and A. Farhadi. Xnor-net: Imagenet classification using binary convolutional neural networks. In *ECCV*, 2016.
- [39] O. Russakovsky, J. Deng, H. Su, J. Krause, S. Satheesh, S. Ma, Z. Huang, A. Karpathy, A. Khosla, M. Bernstein, et al. Imagenet large scale visual recognition challenge. *International Journal of Computer Vision*, 115(3):211–252, 2015.
- [40] T. Salimans and D. P. Kingma. Weight normalization: A simple reparameterization to accelerate training of deep neural networks. In *Advances in Neural Information Processing Systems*, pages 901–909, 2016.
- [41] 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*, pages 4510–4520, 2018.
- [42] W. Shang, J. Chiu, and K. Sohn. Exploring normalization in deep residual networks with concatenated rectified linear units. In *AAAI*, pages 1509–1516, 2017.
- [43] D. Soudry and Y. Carmon. No bad local minima: Data independent training error guarantees for multilayer neural networks. *arXiv preprint arXiv:1605.08361*, 2016.
- [44] STMicroelectronics. CMOS065LPGP Standard Cell Library. User Manual & Databook, 2008.
- [45] I. Sutskever, J. Martens, G. Dahl, and G. Hinton. On the importance of initialization and momentum in deep learning. In *International conference on machine learning*, pages 1139–1147, 2013.
- [46] W. Tang, G. Hua, and L. Wang. How to train a compact binary neural network with high accuracy? In *AAAI*, pages 2625–2631, 2017.
- [47] T. Tieleman and G. Hinton. Lecture 6.5-rmsprop: Divide the gradient by a running average of its recent magnitude. *COURSERA: Neural networks for machine learning*, 4(2):26–31, 2012.
- [48] D. Xie, J. Xiong, and S. Pu. All you need is beyond a good init: Exploring better solution for training extremely deep convolutional neural networks with orthonormality and modulation. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 6176–6185, 2017.
- [49] M. Yazdani. Linear backprop in non-linear networks. 2018.
- [50] S. Zagoruyko and N. Komodakis. Wide residual networks. *arXiv preprint arXiv:1605.07146*, 2016.
- [51] R. Zhao, W. Song, W. Zhang, T. Xing, J.-H. Lin, M. Srivastava, R. Gupta, and Z. Zhang. Accelerating binarized convolutional neural networks with software-programmable fpgas. In *Proceedings of the 2017 ACM/SIGDA International Symposium on Field-Programmable Gate Arrays*, pages 15–24. ACM, 2017.
- [52] 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.