

Article

IE-Net: Information-Enhanced Binary Neural Networks for Accurate Classification

Rui Ding, Haijun Liu *  and Xichuan Zhou

School of Microelectronics and Communication Engineering, Chongqing University, Chongqing 400044, China; dingrui961210@cqu.edu.cn (R.D.); zxc@cqu.edu.cn (X.Z.)

* Correspondence: haijun_liu@126.com

Abstract: Binary neural networks (BNNs) have been proposed to reduce the heavy memory and computation burdens in deep neural networks. However, the binarized weights and activations in BNNs cause huge information loss, which leads to a severe accuracy decrease, and hinders the real-world applications of BNNs. To solve this problem, in this paper, we propose the information-enhanced network (IE-Net) to improve the performance of BNNs. Firstly, we design an information-enhanced binary convolution (IE-BC), which enriches the information of binary activations and boosts the representational power of the binary convolution. Secondly, we propose an information-enhanced estimator (IEE) to gradually approximate the sign function, which not only reduces the information loss caused by quantization error, but also retains the information of binary weights. Furthermore, by reducing the information loss in binary representations, the novel binary convolution and estimator gain large information compared with the previous work. The experimental results show that the IE-Net achieves accuracies of 88.5% (ResNet-20) and 61.4% (ResNet-18) on CIFAR-10 and ImageNet datasets respectively, which outperforms other SOTA methods. In conclusion, the performance of BNNs could be improved significantly with information enhancement on both weights and activations.

Keywords: binary neural networks; deep learning; information enhancement; image classification



Citation: Ding, R.; Liu, H.; Zhou, X. IE-Net: Information-Enhanced Binary Neural Networks for Accurate Classification. *Electronics* **2022**, *11*, 937. <https://doi.org/10.3390/electronics11060937>

Academic Editor: Amir Mosavi

Received: 1 March 2022

Accepted: 16 March 2022

Published: 17 March 2022

Publisher's Note: MDPI stays neutral with regard to jurisdictional claims in published maps and institutional affiliations.



Copyright: © 2022 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

1. Introduction

In recent years, due to great representational power and the good ability to process image data, deep convolutional neural networks (DCNNs) have been used in various computer vision tasks, such as image classification [1,2], object detection [3,4], and semantic segmentation [5,6]. It is reported that most of the modern, powerful DCNNs need a considerable number of learnable parameters and computation requirements, which impose high demands on the hardware that supports their running. However, with the advent of the Internet of Things, how to deploy high-performance DCNNs on embedded devices with limited hardware resources has become an urgent problem. To solve this problem, many model compression methods which reduce the model size and computational burden have been proposed, such as network quantization [7,8], model pruning [9], knowledge distillation [10,11], and lightweight model design [12].

Among them, network quantization is regarded as a simple yet effective solution, where the activations and weights are represented by the lower bits. Binary neural networks (BNNs) [13] are the extreme versions of the quantized neural networks, which binarize both the activations and weights within the network to discrete values $\{+1, -1\}$. Through this method, one could store the model parameters with only 1-bit representations. Moreover, the floating-point operations (FLOPs) in DCNNs can be replaced with the cheap logic operations (XNOR and POPCOUNT), due to the 1-bit advantage. In summary, BNNs could reduce the memory and computation requirements of DCNNs significantly, which shows a great potential to solve the problem of model deployment on embedded devices.

However, the BNN methods often induce a large performance drop compared with the full-precision counterparts. For example, directly applying the normal binary technology [13] on the AlexNet model causes 28.7% Top-1 accuracy loss on the ImageNet dataset [14]. The reason for the accuracy decrease is that the 1-bit representation of the activations and weights reduces the representational power of the BNNs, and leads to huge information loss during both training and inference time. To tackle the problems, a lot of related works have been proposed. IR-Net [15] proposes to use the balance and standardization operations before binarizing the weights to maximize the information entropy of weights and activations. ReActNet [16] adopts the sign function with learnable thresholds (RSign) to binarize the activations, which reduces the information loss of the activation features. Furthermore, ABC-Net [17], BENN [18], GroupNet [19], and CBCN [20] improve the representational power and increase the information of the BNNs by adopting more binary bases, which leads to extra memory and computation costs. Although the above methods alleviate the information loss and enlarge the representational power, the information of the binary models could be enhanced even further. Firstly, the ReActNet only uses a single RSign function to generate the binary activations, which may lose some useful and diverse information from original feature maps. Secondly, the approximated function error decay estimator (EDE) proposed by IR-Net does not provide a strong gradient signal enough to help the weights decide their signs, which leads to the suboptimal results of the information increase and the quantization error reduction. Finally, due to the additional binary bases, the methods like BENN harm the hardware-friendly properties of BNNs.

In this work, we propose to build the information enhanced network (IE-Net) with binary weights and activations. Figure 1 shows the forward and backward processes of the proposed binary neural network.

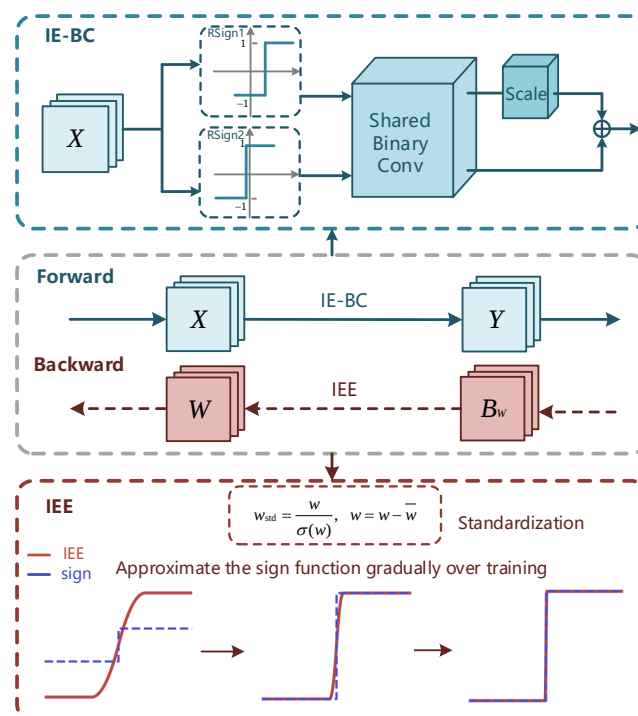


Figure 1. Overview of the proposed IE-Net. The information enhanced binary convolution (IE-BC) is used to enrich the information of binary activations and boost the representation power of binary convolution. The information enhanced estimator (IEE) is proposed to help maximize the information entropy of weights and minimize the quantization error to reduce the information loss.

In the forward process, the information-enhanced binary convolution (IE-BC) has been proposed to improve the representational power of the binary convolutional layers. The IE-BC adopts multiple sign functions with learnable thresholds to enrich the information

of the binary activations. Besides, the multiple binarized activations are processed by the shared binary convolution. To generate more diverse features from the same convolution, we employ scaling factors on the output feature maps, which adds minor memory and computation costs. In the backward process, we propose the information-enhanced estimator (IEE) to optimize the standardized balanced weights, which helps latent weights decide their signs to minimize the quantization error and maximize the information entropy of the binary weights. With the help of the proposed IE-Net, one could train an accurate binary neural network with information enhancement.

In summary, the main contributions of the proposed method are threefold:

1. To enhance the information of binary activations and improve the representational power of the binary model, we present the novel binary function IE-BC, which employs multiple sign functions with different learnable thresholds, a shared binary convolution, and following scaling operations. The diverse binary activations generated by IE-BC retain the information of the original input, and the novel convolution in IE-BC could combine the multiple binary features effectively with information enhancement. In addition, the IE-BC improves the model performance with minor memory and computation requirements increase.
2. To help weights decide their signs and achieve better information gain on binary weights, we propose to replace the STE method with the IEE function that approximates the original sign function as training proceeds. With the help of the proposed method, the gradients of the weights are adapted according to different training stages and strong enough to help the weights update. The IEE could shape the weights distributions around +1 and −1, which reduces the quantization error that introduces information loss and maximizes the information entropy of weights in each layer.
3. The experimental results show that our proposed IE-Net increases the mean accuracy of the baseline model Bi-Real [21] by 2.8% and outperforms the other state-of-the-art (SOTA) BNN methods on the CIFAR-10 dataset. Besides, we evaluate our method with the ResNet-18 and ResNet-34 [1] structures on the ImageNet dataset and the results show that the IE-Net achieves the best performance compared with other SOTA models, which proves the effectiveness of the proposed method.

2. Materials and Methods

2.1. Binary Neural Networks

The BNN [13] firstly proposes the method to binarize the activations and weights in deep neural networks and introduces the strategy for training the BNNs. This kind of model compression technology could reduce the memory and computation requirements significantly. Technically, for most recent BNN methods, both the weights and activation inputs in the convolutional layers are binarized with a sign function, and their specific formulas are shown as follows:

$$b_w = \text{sign}(w) = \begin{cases} +1 & \text{if } w \geq 0 \\ -1 & \text{otherwise} \end{cases}, \quad b_x = \text{sign}(x) = \begin{cases} +1 & \text{if } x \geq 0 \\ -1 & \text{otherwise} \end{cases} \quad (1)$$

where the x and w are the elements of full-precision activations X and full-precision weights W , and b_x and b_w are the elements of binary activations B_X and binary weights B_W respectively. According to Equation (1), the weights and activations could be binarized to $\{+1, -1\}$, which saves the memory storage by $32\times$ in theory. By taking advantage of the binary representations, the energy-hungry floating-point operations in the original full-precision convolution are replaced with efficient logical operations.

$$Y = (B_W \otimes B_X) \cdot \alpha \quad (2)$$

where the $\otimes$ is denoted as a bitwise operation XNOR and POPCOUNT and the α is a scaling factor that is used to reduce the quantization error caused by the binarization. In recent years, a lot of BNN methods have been proposed to improve the performance of binary

models by minimizing the quantization error in either magnitude [14,21] or angular [22,23] aspects. Finally, in the forward inference time, the binary neural networks could save the memory and computation costs significantly.

During the training process, the parameters in BNNs are hard to update since the gradients of the sign function are nearly zero almost everywhere. To solve this problem, the straight through estimator (STE) [24] has been adopted to back-propagate the gradient through the sign function as follows:

$$\frac{\partial \mathcal{L}}{\partial W} = \frac{\partial \mathcal{L}}{\partial B_W} \frac{\partial B_W}{\partial W} \approx \text{clip}\left(\frac{\partial \mathcal{L}}{\partial B_W}, -1, +1\right), \quad (3)$$

where the clip function is a piece-wise linear function, $\mathcal{L}$ is the loss function of the binary neural network and W is used as the latent weights to be updated in the backward process. However, the approximated gradients using STE lead to the gradient mismatch problem, which influences the training of BNNs. Previous methods [15,21,22] propose to use alternative functions which approximate the sign function to reduce this mismatch problem and help train the BNNs.

2.2. Information Enhanced Binary Convolution (IE-BC)

Due to the very limited representations of weights and activations, the normal BNNs show a large performance degeneration compared with the full-precision counterparts. In particular, activations are more sensitive to the binarization process. BinaryConnect [25] only binarizes the weights in DNNs and finds that the final model with 1-bit weights and 32-bit activations achieves comparable results with the full-precision model on small-scale datasets due to the regularization effect of binary weights. Furthermore, the experiments of IR-Net [15] also show that the models that only binarize the weights and keep activations in full-precision improve the accuracy significantly compared with the normal BNNs. Therefore, directly binarizing activations may bring a large information loss, which decreases the final accuracy of binary models. To solve this problem, we propose the information enhanced binary convolution IE-BC to reduce the information loss in activations.

In the forward process of most BNN methods, the input activations of the binary convolution are binarized to $\{+1, -1\}$ with a sign function. Due to the limited values in activations, the input features become redundant and uninformative. To reduce the degradation of activations, ReActNet [16] proposes the RSign function which is defined as a sign function with channel-wise learnable thresholds. However, as shown in Figure 2, the thresholds within each channel may not find the appropriate value, and the binary activation generated by the sign function with a bad threshold leads to huge information loss which harms the model performance. Thus, we propose to adopt the multiple RSign functions with different channel-wise learnable thresholds, which is formulated as follows:

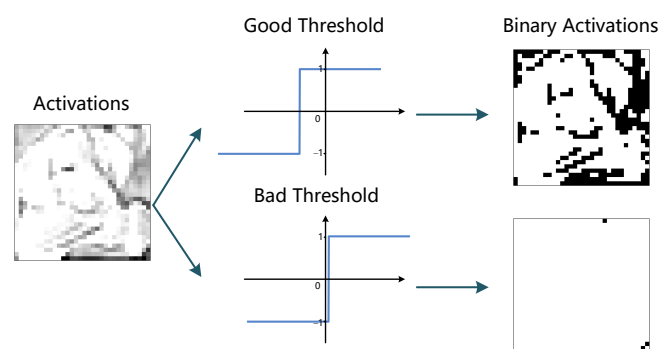


Figure 2. The visualization of the binary activations generated by different sign functions. The sign function above with a negative threshold generates the informative binary activation which retains the information of the original input. The sign function below with a small positive threshold generates the meaningless binary activation which causes huge information loss.

$$b_x^{i,k} = h^k(x^i) = \begin{cases} +1 & \text{if } x^i \geq \beta^{i,k} \\ -1 & \text{if } x^i < \beta^{i,k} \end{cases} \quad (4)$$

where the $b_x^{i,k}$ is the i th channel binary activation element generated by the k th sign function h^k with its i th channel learnable threshold $\beta^{i,k}$, and we denote K as the total number of the RSign function used in IE-BC. According to Equation (4), we derive multiple groups of binary activation inputs. After that, a simple way is to use different binary convolutions on different generated binary activations, respectively. However, this method causes the linear growth of memory and computation burdens when the K goes larger. To alleviate the extra complexity, we propose to use a shared binary convolution to deal with all these binary activations as follows:

$$Y^k = \text{BConv}(B_W, B_X^k) = (B_W \otimes B_X^k) \cdot \alpha \quad (5)$$

where the B_X^k is the k th binary activation, Y^k is the k th output and B_W is the shared convolutional weights. Although the shared convolution saves the memory and computation costs, the same convolutional filters could harm the diversity of the output features, which influences the representational power of the binary model. To tackle this challenge, we propose to apply the channel-wise scaling factors to compensate for the diversity loss and enrich the information that outputs contain. Then the final output Y is computed as follows:

$$Y = Y^1 + \sum_{k=2}^K Y^k \cdot \lambda^k \quad (6)$$

where the λ^k is the compensation factor that scales the k th output generated by the shared convolution with the k th binary activation. It is worth noting that the first output is considered as the base output activation which has no need to use the compensation factor. In addition, the experiment section will analyze the influence of the hyperparameter K and the results show that the $K = 2$ is the best choice which achieves the highest accuracy with little complexity increase.

In summary, the whole structure of the IE-BC ($K = 2$) is illustrated in Figure 3. As the figure shows, the sign functions with two different thresholds generate two binary activations which contain completely different information. With this novel binary convolution, we could enhance the information of binary activations and improve the representational power of binary neural networks.

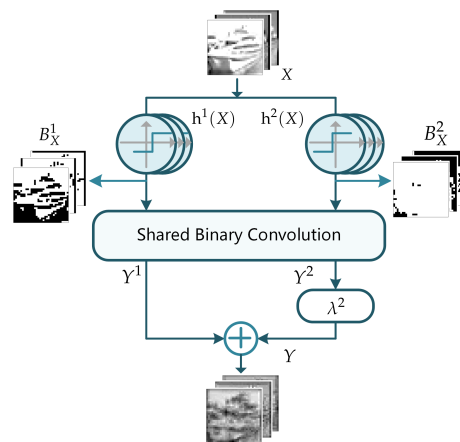


Figure 3. The structure of the IE-BC when $K = 2$. The 32×32 input activation X is binarized to the different binary activations B_X^1 and B_X^2 by two RSign functions. Then, a shared binary convolution is used to process the binary activations and generates the outputs Y^1 and Y^2 . In the end, the final output is derived by the summation of Y^1 and $Y^2 \cdot \lambda^2$, according to Equation (6).

2.3. Information-Enhanced Estimator (IEE)

In binary neural networks, the process of binarization always introduces the large quantization error, which leads to huge information loss. To reduce the loss, many binary works [7,14,21,26] have been proposed to minimize the quantization error with different optimized methods. Besides, IR-Net [15] proposes that only aiming to narrow the difference between full-precision and binary weights will harm the information entropy of binary weights, which hurts the training performance. For maximizing the information entropy of the weights in BNNs, IR-Net proposes to balance and standardize the weights, before binarizing them as follows:

$$w_{\text{std}} = \frac{\hat{w}}{\sigma(\hat{w})}, \quad \hat{w} = w - \bar{w} \quad (7)$$

where $\bar{w}$ denotes the mean value and $\sigma(\cdot)$ means the standard deviation.

As shown in Figure 4, different from the Bi-Real [21], the IR-Net optimizes the weights to form a bimodal distribution which increases the information entropy. However, from the figure, we could find that the two peaks of the weight distribution in IR-Net are not centered on +1 and −1, and there are still some of the weights around the zero value. Thus, the quantization error is relatively large which causes additional information loss that influences the final performance.

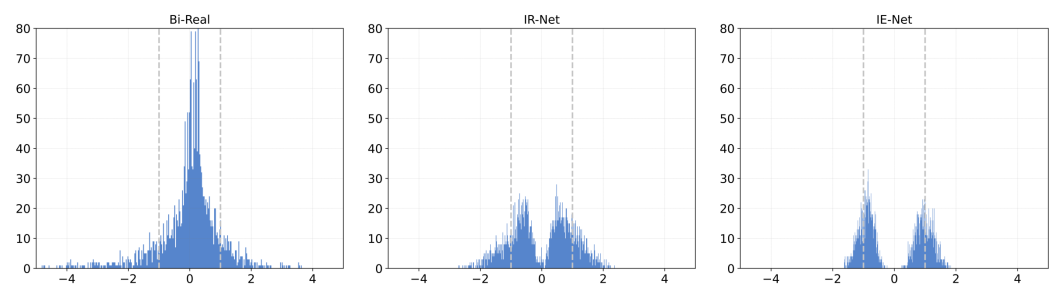


Figure 4. Weight distributions (before binarization) of the Bi-Real, IR-Net, and IE-Net in layer1.0.conv1 of binary ResNet-20. The weights in Bi-Real are gathered around zero value, which is far from the binary values +1 and −1, leading to large quantization error. The IR-Net alleviates this problem while there are still a relatively large number of weights existing around zero. The IE-Net proposes the IEE to shape the two distributions around +1 and −1, which enhances the information of binary weights by reducing the quantization error.

To solve this problem, we propose to combine the idea from IR-Net with a novel training-aware estimator and optimizes the weights to be distributed like the third subfigure in Figure 4, which enhances the information of binary weights. For maximizing the information entropy, we balance and standardize the full-precision weights according to Equation (7) before the binarization, which modifies the weight distributions. For minimizing the quantization error, we propose that the IEE could give the weights strong gradients to help them decide their signs and gradually approximate the sign function to reduce the gradient mismatch problem. Additionally, to make the model training more stable, various approximated functions have been proposed to help update the parameters in BNNs, such as STE [24], piecewise polynomial function [21], EDE [15], and so on.

In this paper, without using existing estimators, we propose a gradually adapted function IEE with the training processes which formula is shown as follows:

$$F(x) = \begin{cases} r(-\text{sign}(x)\frac{3q^2x^2}{4} + \sqrt{3}qx) & \text{if } |x| < \frac{2\sqrt{3}}{3q} \\ r\text{sign}(x) & \text{otherwise} \end{cases} \quad (8)$$

where the r and q are the variables that control the shape of the IEE during the training:

$$q = 10^{T_{\min} + \frac{e}{E}(T_{\max} - T_{\min})}, \quad r = \max\left(\frac{1}{q}, 1\right) \quad (9)$$

where we set $T_{\min} = -2$, $T_{\max} = 1$ in this work, e and E denote the current training epoch and the total number of epochs respectively. According to the Equations (8) and (9), the $F(x)$ could gradually approximate the sign function with the training process which is indicated by the value of $\frac{e}{E}$. In the backward pass, the gradient of IEE concerning the input x could be computed by the following formula:

$$F'(x) = \frac{\partial F(x)}{\partial x} = \begin{cases} r(\sqrt{3}q + \frac{3q^2x}{2}) & \text{if } -\frac{2\sqrt{3}}{3q} \leq x < 0 \\ r(\sqrt{3}q - \frac{3q^2x}{2}) & \text{if } 0 \leq x < \frac{2\sqrt{3}}{3q} \\ 0 & \text{otherwise.} \end{cases} \quad (10)$$

Then, we could derive the gradients of the loss function $\mathcal{L}$ with respect to weights W :

$$\frac{\partial \mathcal{L}}{\partial W} = \frac{\partial \mathcal{L}}{\partial B_W} F'(W) \quad (11)$$

Besides, to intuitively understand the proposed IEE, we visualize the function shape of $F(x)$ and $F'(x)$ with growing value of $\frac{e}{E}$ in Figure 5. As the figure demonstrates, at the beginning of the training phase, the gradients exist almost everywhere and have a larger value than 1 compared with the other estimators [21,24], which encourages the weights to flip their signs and help optimize the binary model. As the training goes on, the shape of the $F'(x)$ gradually fits with the gradients of the sign function, which reduces the gradient mismatch problem. Furthermore, the magnitude of the gradients becomes even larger, which helps the weights decide their signs. After training the model, the weights in BNNs are pushed to gather around +1 and −1, resulting in the minimized quantization error. Meanwhile, due to the balance and standardization operations before binarization, the information entropy is also optimized at the same time.

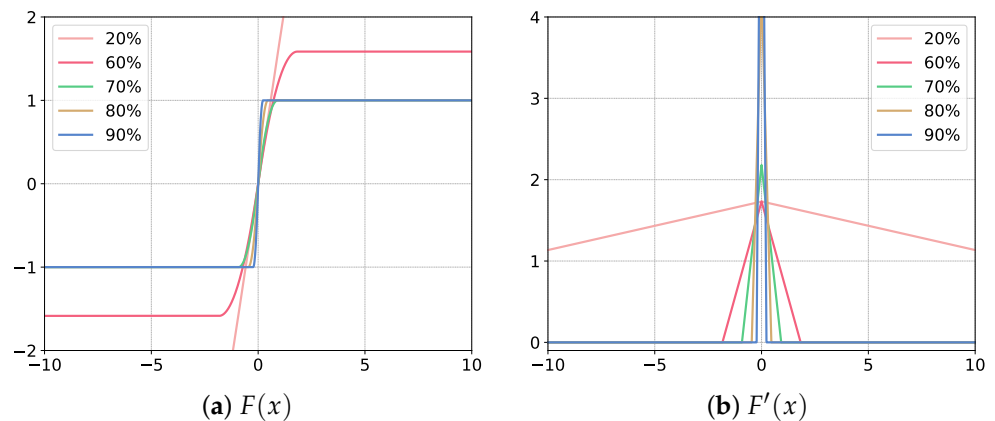


Figure 5. Visualization of the IEE in the different training stages which are indicated by $\frac{e}{E} \times 100\%$. (a) The function shape of $F(x)$. (b) The derivative shape of $F'(x)$.

3. Experiments and Discussion

3.1. Experimental Settings

3.1.1. Datasets

CIFAR-10: CIFAR-10 [27] is a computer vision dataset collected by the students of the Hinton group for pervasive object recognition, which contains 10 categories. The dataset consists of 60,000 color images of size 32×32 , of which 50,000 images are used for training the models and 10,000 images are used for evaluating the model performance. Like most

previous works, the data augmentation methods including random crop and flipping are adopted during the training phase, while dataset normalization is used in both the training and testing phases.

ImageNet: ILSVRC 2012 ImageNet [28] is a large-scale and high-resolution image dataset which contains 1000 classes for image recognition. The dataset includes 1.2 million natural RGB images used for training and 50,000 RGB images for evaluation. The commonly used data augmentation strategies such as random crop and random flipping are adopted during the training process. In the testing phase, we evaluate our models on 224×224 center-cropped images from the testing set.

3.1.2. Implementation Details

All the experiments in this section are implemented with the powerful and flexible Pytorch library and conducted on a single computer with an Intel Xeon E5-2680 CPU and 4 NVIDIA RTX 3090 GPUs. Following the compared binary neural networks, we binarize all the convolutional layers and fully connected layers, except the first and last layers. Besides, our source code is accessed on 17 March 2022 at <https://github.com/Alexrich961210/IE-Net>.

For the experiments on the CIFAR-10 dataset, we use the Bi-Real method based on ResNet-20 as the baseline model to conduct the ablation studies. Besides, we also evaluate our method on VGG-Small and ResNet-18 network topologies, respectively. In the training time, we choose the SGD optimizer with the momentum of 0.9 as the default optimizer. The weight decay is set as 1×10^{-4} and the batch size is 128. The initial learning rate is set as 0.1 and we adopt the cosine annealing strategy to adjust the learning rate as the training processes. In addition, we train all the models 400 epochs in total.

For the experiments on the ImageNet dataset, we evaluate our method based on ResNet-18 and ResNet-34 network structure, respectively. During the training process, the SGD optimizer with the momentum of 0.9 is adopted, the weight decay is set as 1×10^{-4} , and the batch size is set as 512. The initial learning rate is set as 0.1. and the cosine annealing scheduler is used to adjust the learning rate. The warm-up strategy is also used to help the binary models converge. All the models are trained for 120 epochs.

3.2. Ablation Study

In this part, we explore and analyze the effects of the proposed IE-BC and IEE on binary neural networks. We use the ResNet-20 based Bi-Real as the baseline model which applies double skip connections and we evaluate all the models on the CIFAR-10 dataset. Additionally, we set the hyperparameter K which denotes the number of used modified sign functions as 2 by default.

3.2.1. Effectiveness of Information Enhanced Binary Convolution (IE-BC)

Based on the Bi-Real Net with ResNet-20 network, we replace the binary convolution within the baseline model with the proposed enhanced binary convolution (EBC) to test its influence on the model performance. Besides, to verify the advantages of the proposed EBC method on the representational power improvement, we compare our method with the RSign technology proposed by ReActNet [16].

Table 1 lists the classification accuracy of the Bi-Real baseline, the Bi-Real+RSign, and the Bi-Real+IE-BC on CIFAR-10, respectively. All of them use the same training settings and network structures. RSign inserts a learnable shift parameter before the sign function to improve the quality of the binary feature maps. The IE-BC applies the sign functions with multiple learnable thresholds to generate diverse binary patterns. From the table, it is clear to see that our proposed IE-BC method improves the performance of the baseline model significantly, increasing the mean accuracy by 2.60%. In addition, the proposed Bi-Real+IE-BC outperforms the Bi-Real+RSign method by a mean accuracy of 1.46%, which proves the superiority of the novel binary convolution method.

Table 1. Comparison of the different methods that use modified sign function on the baseline model. We run each model three times and report the mean accuracy and standard deviation on the CIFAR-10 test dataset.

	Bi-Real Baseline [21]	Bi-Real+RSign [16]	Bi-Real+IE-BC
Mean Accuracy (%)	85.74	86.88	88.34
Std (%)	0.19	0.27	0.15

Furthermore, to display the reason why the IE-BC brings a large performance gain intuitively, we visualize the feature maps within the first binary convolutional layer. Figure 6 shows the binary activation inputs at the first 4 channels after multiple sign functions with different learned thresholds inserted in the IE-BC module. As the figure demonstrates, we can find that:

1. For the same full-precision inputs, the binarized inputs generated by the sign functions with different thresholds present diverse features. In particular, the information of binary inputs in the 1th channel is completely different, which helps the binary model learn more meaningful patterns.

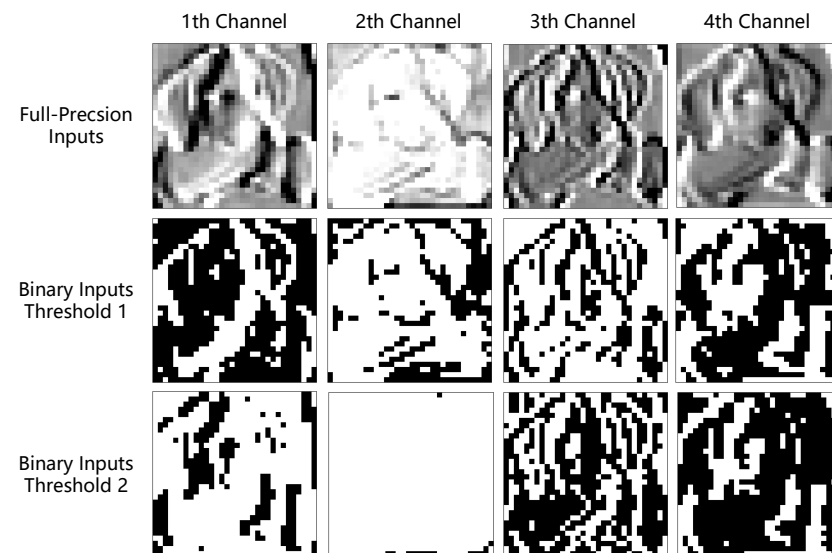


Figure 6. Visualization results of the feature maps in the first binary convolutional layer of Bi-Real+IE-BC based on ResNet-20. The rows indicate the full-precision inputs and binary inputs with different learned thresholds. The columns indicate the feature maps from the first four channels. The quality of the binary activations is measured by the visual difference between the same full-precision input and different binarized inputs.

2. Meanwhile, as feature maps from the 2th channel show, the sign function with a bad threshold will generate meaningless binary activations as shown in the third row which induces large information loss. By using the IE-BC method, the binary activations from another sign function with a different threshold could compensate the missing feature as shown in the second row at 2th channel, which proves the effectiveness of the proposed technology.
3. In conclusion, the different activation binarized functions could generate multiple diversified binary patterns to help enhance the information of binary activations and boost the representational power of the normal binary convolution, which increases the final classification accuracy of the binary models.

3.2.2. Influence of Hyperparameter K

Considering the great improvement from using 2 modified sign functions with their thresholds in the IE-BC, a natural question is whether more modified sign functions lead to better model performance. We denote the number of the used modified sign function as K , which is consistent with the last method section.

To explore this question, we conduct a group of experiments using Bi-Real+IE-BC with K tuned from 1 to 5, and the experimental results are shown in Figure 7. It could be seen that the performance improvement on the baseline model is significant when K changes from 1 to 2, and gradually becomes smaller when K is greater than 2. The lowest mean error rate 11.66% is achieved when $K = 2$. The experimental phenomenon means that a more number of the modified sign functions may be redundant for the good performance of the binary model. Besides, the redundant sign functions will introduce more memory cost and computational complexity due to the learnable thresholds, which influences the hardware-friendly nature of the binary neural networks to some extent. Therefore, we choose K as 2 to build the final binary neural network, which improves the performance greatly and introduces minor memory and computation burdens.

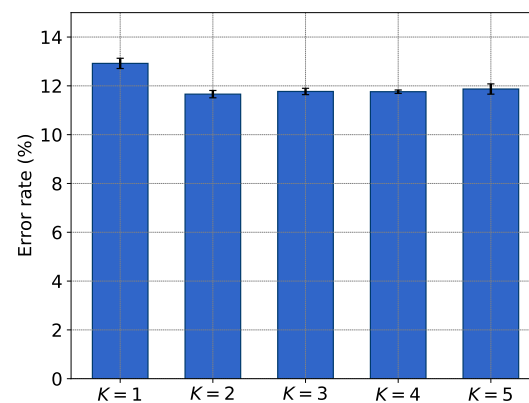


Figure 7. The error rate of Bi-Real+IE-BC with respect to a different number K of modified sign functions. The error rate is obtained by the difference between 1 and the accuracy rate on the CIFAR-10 dataset, and the lower error rate indicates better model performance.

3.2.3. Effectiveness of Information-Enhanced Estimator (IEE)

To show the effect of the proposed IEE function on minimizing the quantization error, we visualize the data distribution of full-precision weights and the derivative of the proposed approximated function during the training process, where the results are shown in Figure 8.

At the early stage of the training (10 epoch and 200 epoch), there exist many full-precision weights outside the range from -1 to $+1$, which can not be updated using STE according to Equation (3). To solve this problem, IEE relaxes the truncation threshold to let the gradients exist for almost all the weights and enlarge the magnitude of the gradients to help the weights flip their signs which is beneficial for training the binary model. Besides, it could be seen that the derivative curve of the IEE becomes more similar to the sign function as the training processes, which reduces the gradient mismatch problem. Meanwhile, in the last stage of the training (400 epoch), the weights finally gather around the value of $+1$ and -1 and shape a clear gap between these two data distributions. As the figure shows, there are no weights around zero value, which reduces the quantization error effectively and thus obtain the information gain.

To verify the effectiveness of the proposed IEE on maximizing the information entropy, we compare the total information entropy of the weights with a Bi-Real baseline model, IR-Net [15] and Bi-Real+Median Loss (ML) [29] based on the ResNet-20, as shown in Table 2.

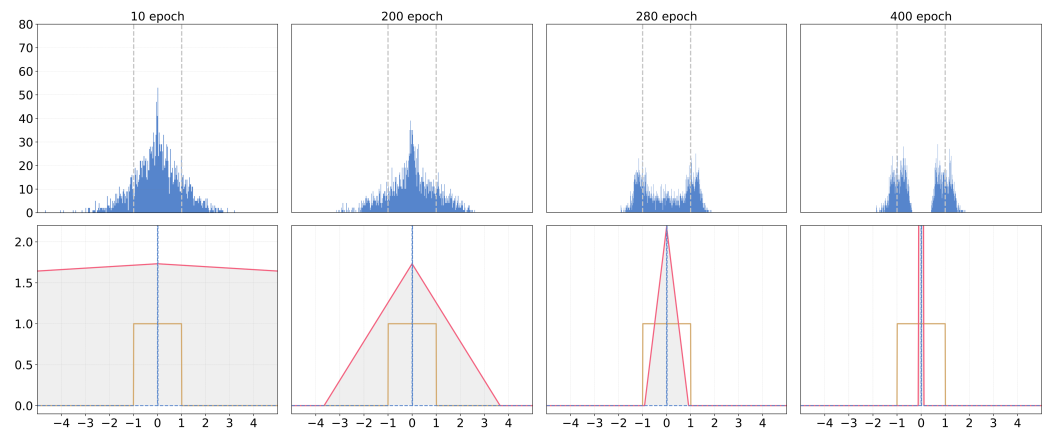


Figure 8. Visualization of the weight (after balance and standardization) distributions and derivatives of IEE in different training epochs (10, 200, 280, and 400). The first row demonstrates the distributions, where the gray vertical lines indicate the -1 and $+1$ values. The second row demonstrates the corresponding derivative curves, where the red curve is the proposed IEE, the yellow curve is the STE and the blue one is the sign function.

As the table demonstrates, the IEE increases the information entropy of the baseline model from 5.39 to 5.42, which proves the effect of IEE on enhancing the information of binary weights. In addition, the Bi-Real+IEE achieves the equivalent result compared with the IR-Net. However, the IEE could enjoy the information gain by minimizing the quantization error at the same time, which helps improve the model performance, especially in large-scale datasets.

Table 2. Comparison with other related works on increasing information entropy of binary weights. We run all the models three times and report the mean information entropy on the CIFAR-10 test dataset. The evaluation metric is defined as the summation of the information entropy of binary weights in all the binary convolutional layers.

Metric	Bi-Real [21]	IR-Net [15]	BiReal+ML [29]	Bi-Real+IEE
$\sum_{l=2}^{19} \mathcal{H}(l)$	5.39	5.42	5.41	5.42

3.2.4. Ablation Performance

In this part, we demonstrate the effect of different elements in the proposed IE-Net on the performance of the binary baseline model with ResNet-20 structure, where the experimental results on the CIFAR-10 dataset are shown in Table 3. From the Table, it is reported that the IE-BC and IEE modules increase the mean accuracy of the baseline model by 2.60% and 0.73% respectively, which proves the effectiveness of the proposed methods. Furthermore, by combining these two components together, we derive the proposed IE-Net and achieve even better results, increasing the mean accuracy by 2.80% compared with the baseline model, which proves that enhancing the information of binary activations and weights could improve the performance of the binary model significantly.

Table 3. Ablation performance on the CIFAR-10 dataset. We run each model three times and report the mean $\pm$ std accuracy on the test dataset.

Topology	Method	Bit-Width (W/A)	Accuracy (%)
ResNet-20	Bi-Real	1/1	85.74 ± 0.19
	+IE-BC	1/1	88.34 ± 0.15
	+IEE	1/1	86.47 ± 0.09
	+IE-BC+IEE (IE-Net)	1/1	88.54 ± 0.14

3.3. Comparison with State-of-the-Art Methods

In this section, we comprehensively compare the proposed IE-Net with other SOTA methods on CIFAR-10 and ImageNet datasets respectively.

3.3.1. Comparisons on CIFAR-10

Table 4 compares the classification accuracy of the IE-Net with other SOTA binary neural networks on the CIFAR-10 dataset, including RAD [30], IR-Net [15] and RBNN [22] based on ResNet-18; DoReFa [31], DSQ [32], XNOR+ML+BMA [29], SLB [33], IR-Net and RBNN based on ResNet-20; XNOR-Net [14], BNN [13], IR-Net, RAD, RBNN and DSQ based on VGG-Small. As the table shows, our proposed binary model achieves the best performance in different network structures, which verifies the universality and superiority of our method. It is worth noting that the IE-Net yields 1.4%, 2%, and 1.6% performance gains over IR-Net based on ResNet-18, ResNet-20, and VGG-Small respectively due to the information enhancement. Furthermore, the IE-Net based on ResNet-18 narrows the accuracy gap between the binary model and the full-precision counterpart to only 1.9%. Last but not the least, the IE-Net based on VGG-Small also reduces the performance gap to 2.1%.

Table 4. Accuracy comparison with the SOTA methods on the CIFAR-10 dataset. We evaluate our proposed IE-Net based on ResNet-18, ResNet-20, and VGG-Small. The proposed networks are highlighted in bold.

Topology	Method	Bit-Width (W/A)	Accuracy (%)
ResNet-18	Full-Precision	32/32	94.8
	RAD	1/1	90.5
	IR-Net	1/1	91.5
	RBNN	1/1	92.2
	Ours	1/1	92.9
ResNet-20	Full-Precision	32/32	92.1
	DoReFa	1/1	79.3
	DSQ	1/1	84.1
	XNOR+ML+BMA	1/1	85.00
	SLB	1/1	85.5
	IR-Net	1/1	86.5
	RBNN	1/1	87.8
	Ours	1/1	88.5
VGG-Small	Full-Precision	32/32	94.1
	XNOR-Net	1/1	89.8
	BNN	1/1	89.9
	IR-Net	1/1	90.4
	RAD	1/1	90.4
	RBNN	1/1	91.3
	DSQ	1/1	91.7
	Ours	1/1	92.0

3.3.2. Comparisons on ImageNet

We further evaluate the performance of the proposed IE-Net on the large-scale ImageNet dataset. Table 5 compares the Top-1 accuracy and Top-5 accuracy of the IE-Net with other SOTA methods, such as XNOR-Net, DoReFa, TBN [34], Bi-Real [21], PDNN [35], IR-Net, BONN [36] and RBNN based on ResNet-18; ABC-Net [17], Bi-Real, IR-Net and RBNN based on ResNet-34. As can be observed in this table, the proposed IE-Net achieves the best performance compared with the other SOTA binary neural networks with both ResNet-18 and ResNet-34 structures. In addition, the IE-Net obtains a better result than the networks with higher-precision representations such as DoReFa and TBN. In the end, the IE-Net shows the better information gain over IR-Net, which improves the performance of the IR-Net by 3.3% and 1.7% Top-1 accuracy based on ResNet-18 and ResNet-34 respectively.

In conclusion, Tables 4 and 5 prove the effectiveness of the proposed method on enhancing the information of binary neural networks, which could improve the final model performance in various network structures.

Table 5. Accuracy comparison with the SOTA methods on the ImageNet dataset. We evaluate our proposed IE-Net based on ResNet-18 and ResNet-34. The proposed networks are highlighted in bold.

Topology	Method	Bit-Width (W/A)	Top-1 (%)	Top-5 (%)
ResNet-18	Full-Precision	32/32	69.6	89.2
	XNOR-Net	1/1	51.2	73.2
	DoReFa	1/2	53.4	-
	TBN	1/2	55.6	79.0
	Bi-Real	1/1	56.4	79.5
	PDNN	1/1	57.3	80.0
	IR-Net	1/1	58.1	80.0
	BONN	1/1	59.3	81.6
	RBNN	1/1	59.9	81.9
	Ours	1/1	61.4	83.0
ResNet-34	Full-Precision	32/32	73.3	91.3
	ABC-Net	1/1	52.4	76.5
	Bi-Real	1/1	62.2	83.9
	IR-Net	1/1	62.9	84.1
	RBNN	1/1	63.1	84.4
	Ours	1/1	64.6	85.2

3.4. Memory and Computation Complexity Analyses

In this part, we analyze the model complexity of the proposed IE-Net based on ResNet-18 and ResNet-34. Table 6 demonstrates the memory computational costs of different binary neural networks, including XNOR-Net, Bi-Real, IR-Net, and the proposed IE-Net.

Table 6. Comparison of memory cost and computation complexity with different methods based on ResNet-18 and ResNet-34. The memory cost is represented by the number of bits occupied by the model parameters, and the computational cost is denoted by the floating-point operations within the binary model. The proposed networks are highlighted in bold.

Topology	Method	Bit-Width (W/A)	Memory Cost (Mbit)	FLOPs
ResNet-18	Full-Precision	32/32	374.1	1.81×10^9
	XNOR-Net	1/1	33.7	1.67×10^8
	Bi-Real	1/1	33.6	1.63×10^8
	IR-Net	1/1	33.6	1.63×10^8
	Ours	1/1	33.8	1.63×10^8
ResNet-34	Full-Precision	32/32	697.3	3.66×10^9
	XNOR-Net	1/1	43.9	1.98×10^8
	Bi-Real Net	1/1	43.7	1.93×10^8
	IR-Net	1/1	43.7	1.93×10^8
	Ours	1/1	44.1	1.93×10^8

Compared with the full-precision counterparts, the IE-Net saves the memory and computational costs by $11.07 \times / 15.81 \times$ and $11.10 \times / 18.96 \times$ based on ResNet-18/34 respectively. From the table, it is clear to see that the model complexity reduction of our method is comparable with other listed BNN methods while our model achieves better performance for image classification, according to Table 5.

Furthermore, the extra memory and computation requirements introduced by the IE-BC module are also computed. For the memory cost, the IE-BC in the IE-Net only adopts two modified sign functions to enrich the information of binary activations, which adds the 0.2 Mbit and 0.4 Mbit storage costs based on ResNet-18 and ResNet-34 respectively.

Besides, the compensation factors could be absorbed by the batch norm layers so they will not introduce other complexity in the inference time. For the computational cost, the IE-BC increases the efficient binary convolution operations instead of FLOPs, and this kind of convolution could be implemented in parallel which does not affect the real inference speed. According to Tables 5 and 6, we find that the IE-Net achieves a good trade-off between model complexity and model performance.

In summary, although the IE-Net introduces a few computational and memory costs, the information gain and performance improvement are significant, which benefits the deployment of the binary neural networks on real-world applications.

4. Conclusions

In this work, we propose a novel binary neural network named IE-Net to enhance the information and performance of the binary models. Firstly, we propose an information enhanced binary convolution (IE-BC) to enrich the information of binary activations and boost the representational power of the binary convolution. The IE-BC applies multiple sign functions with multiple learnable thresholds to generate diverse binary input features, which retain more information from original inputs. Then, we employ a shared binary convolution equipped with the compensation factors to derive the final output activations with a little additional model complexity. Secondly, to increase the information of weights at the same time, we proposed the information enhanced estimator (IEE) to gradually approximate the sign function and provide the weights strong gradients to update and decide their signs during the training process. After the training phase, the weight distributions and the information entropy metrics show that the proposed IEE not only reduces the quantization error to alleviate the information loss but also obtains larger information entropy compared with the baseline model. With the help of the information enhancement, the IE-Net based on ResNet-20 achieves an accuracy improvement of 2.8% on the CIFAR-10 dataset compared with the baseline model. Besides, the IE-Net based on ResNet-18 obtains a 5.0% Top-1 accuracy gain on the ImageNet dataset and outperforms the other SOTA methods, which demonstrates the superiority of the proposed method. Moreover, the extra memory and computation costs introduced by the proposed binary model are proved to be relatively small, which shows a good trade-off between model complexity and performance.

By enhancing the information within the binary models, our proposed IE-Net could improve the performance of BNNs significantly. Meanwhile, the performance gains on different network structures also prove the effectiveness of the proposed method. At last, there are two aspects worth to be investigated in the future. Firstly, the model complexity of the IE-Net is able to be further reduced by designing the engineering realization of the IE-BC module, which is beneficial to the deployment of the IE-Net. Secondly, the accurate binary model IE-Net could be used as the backbone network in other computer vision tasks such as object detection and semantic segmentation, which will help the models have a wider range of application scenarios.

Author Contributions: Conceptualization, X.Z.; methodology, R.D.; writing—original draft preparation, R.D. and H.L.; writing—review and editing, H.L. and X.Z. All authors have read and agreed to the published version of the manuscript.

Funding: This work was supported in part by the National Natural Science Foundation of China under Grants 62001063, 61971072 and U2133211, and in part by the China Postdoctoral Science Foundation under Grant 2020M673135, Chongqing Postdoctoral Research Program under Grant XmT2020050.

Conflicts of Interest: The authors declare no conflict of interest.

References

1. He, K.; Zhang, X.; Ren, S.; Sun, J. Deep Residual Learning for Image Recognition. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Las Vegas, NV, USA, 26 June–1 July 2016; pp. 770–778.
2. Wang, X.; Ren, H.; Wang, A. Smish: A Novel Activation Function for Deep Learning Methods. *Electronics* **2022**, *11*, 540. [[CrossRef](#)]

3. Ren, S.; He, K.; Girshick, R.; Sun, J. Faster R-CNN: Towards Real-time Object Detection with Region Proposal Networks. *IEEE Trans. Pattern Anal. Mach. Intell.* **2016**, *39*, 1137–1149. [[CrossRef](#)] [[PubMed](#)]
4. Guo, J.M.; Yang, J.S.; Seshathiri, S.; Wu, H.W. A Light-Weight CNN for Object Detection with Sparse Model and Knowledge Distillation. *Electronics* **2022**, *11*, 575. [[CrossRef](#)]
5. Long, J.; Shelhamer, E.; Darrell, T. Fully Convolutional Networks for Semantic Segmentation. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Boston, MA, USA, 7–12 June 2015; pp. 3431–3440.
6. Xie, X.; Bai, L.; Huang, X. Real-Time LiDAR Point Cloud Semantic Segmentation for Autonomous Driving. *Electronics* **2022**, *11*, 11. [[CrossRef](#)]
7. Zhang, D.; Yang, J.; Ye, D.; Hua, G. Lq-nets: Learned Quantization for Highly Accurate and Compact Deep Neural Networks. In Proceedings of the European Conference on Computer Vision, Munich, Germany, 8–14 September 2018; pp. 365–382.
8. Vandersteegen, M.; Van Beeck, K.; Goedemé, T. Integer-Only CNNs with 4 Bit Weights and Bit-Shift Quantization Scales at Full-Precision Accuracy. *Electronics* **2021**, *10*, 2823. [[CrossRef](#)]
9. Han, S.; Mao, H.; Dally, W.J. Deep Compression: Compressing Deep Neural Networks with Pruning, Trained Quantization and Huffman Coding. In Proceedings of the International Conference on Learning Representations, San Diego, CA, USA, 7–9 May 2015; pp. 1–14.
10. Stewart, R.; Nowlan, A.; Bacchus, P.; Ducasse, Q.; Komendantskaya, E. Optimising Hardware Accelerated Neural Networks with Quantisation and a Knowledge Distillation Evolutionary Algorithm. *Electronics* **2021**, *10*, 396. [[CrossRef](#)]
11. Hinton, G.; Vinyals, O.; Dean, J. Distilling the Knowledge in a Neural Network. In Proceedings of the Advances in Neural Information Processing Systems Workshop, Montreal, QC, Canada, 12–13 December 2014; pp. 1–9.
12. Zhang, X.; Zhou, X.; Lin, M.; Sun, J. Shufflenet: An Extremely Efficient Convolutional Neural Network for Mobile Devices. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Salt Lake City, UT, USA, 19–21 June 2018; pp. 6848–6856.
13. Hubara, I.; Courbariaux, M.; Soudry, D.; El-Yaniv, R.; Bengio, Y. Binarized Neural Networks. In Proceedings of the Advances in Neural Information Processing Systems, Barcelona, Spain, 5–10 December 2016; pp. 4107–4115.
14. Rastegari, M.; Ordonez, V.; Redmon, J.; Farhadi, A. Xnor-net: Imagenet Classification using Binary Convolutional Neural Networks. In Proceedings of the European Conference on Computer Vision, Amsterdam, The Netherlands, 8–16 October 2016; pp. 525–542.
15. Qin, H.; Gong, R.; Liu, X.; Shen, M.; Wei, Z.; Yu, F.; Song, J. Forward and Backward Information Retention for Accurate Binary Neural Networks. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Virtual, 14–19 June 2020; pp. 2250–2259.
16. Liu, Z.; Shen, Z.; Savvides, M.; Cheng, K.T. ReActNet: Towards Precise Binary Neural Network with Generalized Activation Functions. In Proceedings of the European Conference on Computer Vision, Virtual, 23–28 August 2020.
17. Lin, X.; Zhao, C.; Pan, W. Towards Accurate Binary Convolutional Neural Network. In Proceedings of the Advances in Neural Information Processing Systems, Long Beach, CA, USA, 4–9 December 2017; pp. 345–353.
18. Zhu, S.; Dong, X.; Su, H. Binary Ensemble Neural Network: More Bits per Network or More Networks per Bit? In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Long Beach, CA, USA, 16–20 June 2019; pp. 4923–4932.
19. Zhuang, B.; Shen, C.; Tan, M.; Liu, L.; Reid, I. Structured Binary Neural Networks for Accurate Image Classification and Semantic Segmentation. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Long Beach, CA, USA, 16–20 June 2019; pp. 413–422.
20. Liu, C.; Ding, W.; Xia, X.; Zhang, B.; Gu, J.; Liu, J.; Ji, R.; Doermann, D. Circulant Binary Convolutional Networks: Enhancing the Performance of 1-bit Dcnns with Circulant Back Propagation. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Long Beach, CA, USA, 16–20 June 2019; pp. 2691–2699.
21. Liu, Z.; Wu, B.; Luo, W.; Yang, X.; Liu, W.; Cheng, K.T. 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, Munich, Germany, 8–14 September 2018; pp. 722–737.
22. Lin, M.; Ji, R.; Xu, Z.; Zhang, B.; Wang, Y.; Wu, Y.; Huang, F.; Lin, C.W. Rotated Binary Neural Network. In Proceedings of the Advances in Neural Information Processing Systems, Virtual, 6–12 December 2020; pp. 1–12.
23. Xu, S.; Zhao, J.; Lu, J.; Zhang, B.; Han, S.; Doermann, D. Layer-wise Searching for 1-bit Detectors. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Virtual, 19–25 June 2021; pp. 5682–5691.
24. Bengio, Y.; Léonard, N.; Courville, A. Estimating or Propagating Gradients through Stochastic Neurons for Conditional Computation. *arXiv* **2013**, arXiv:1308.3432.
25. Courbariaux, M.; Bengio, Y.; David, J.P. Binaryconnect: Training Deep Neural Networks with Binary Weights during Propagations. In Proceedings of the Advances in Neural Information Processing Systems, Montreal, QC, Canada, 7–12 December 2015.
26. Li, Z.; Ni, B.; Zhang, W.; Yang, X.; Gao, W. Performance Guaranteed Network Acceleration via High-order Residual Quantization. In Proceedings of the IEEE International Conference on Computer Vision, Venice, Italy, 22–29 October 2017; pp. 2584–2592.
27. Krizhevsky, A.; Hinton, G. *Learning Multiple Layers of Features from Tiny Images*; Technical Report; University of Toronto: Toronto, ON, Canada, 2009.
28. Deng, J.; Dong, W.; Socher, R.; Li, L.J.; Li, K.; Li, F. Imagenet: A Large-scale Hierarchical Image Database. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Miami, FL, USA, 20–25 June 2009; pp. 248–255.

29. Zou, W.; Cheng, S.; Wang, L.; Fu, G.; Shang, D.; Zhou, Y.; Zhan, Y. Increasing Information Entropy of Both Weights and Activations for the Binary Neural Networks. *Electronics* **2021**, *10*, 1943. [[CrossRef](#)]
30. Ding, R.; Chin, T.W.; Liu, Z.; Marculescu, D. Regularizing Activation Distribution for Training Binarized Deep Networks. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Long Beach, CA, USA, 15–20 June 2019; pp. 11408–11417.
31. Zhou, S.; Wu, Y.; Ni, Z.; Zhou, X.; Wen, H.; Zou, Y. Dorefa-net: Training Low Bitwidth Convolutional Neural Networks with Low Bitwidth Gradients. *arXiv* **2016**, arXiv:1606.06160.
32. Gong, R.; Liu, X.; Jiang, S.; Li, T.; Hu, P.; Lin, J.; Yu, F.; Yan, J. Differentiable Soft Quantization: Bridging Full-precision and Low-bit Neural Networks. In Proceedings of the IEEE International Conference on Computer Vision, Seoul, Korea, 27 October–2 November 2019; pp. 4852–4861.
33. Yang, Z.; Wang, Y.; Han, K.; Xu, C.; Xu, C.; Tao, D.; Xu, C. Searching for Low-bit Weights in Quantized Neural Networks. In Proceedings of the Advances in Neural Information Processing Systems, Virtual, 7–12 September 2020; pp. 4091–4102.
34. Wan, D.; Shen, F.; Liu, L.; Zhu, F.; Qin, J.; Shao, L.; Shen, H.T. Tbn: Convolutional Neural Network with Ternary Inputs and Binary Weights. In Proceedings of the European Conference on Computer Vision, Munich, Germany, 8–14 September 2018; pp. 315–332.
35. Gu, J.; Li, C.; Zhang, B.; Han, J.; Cao, X.; Liu, J.; Doermann, D. Projection Convolutional Neural Networks for 1-bit Cnns via Discrete Back Propagation. In Proceedings of the AAAI Conference on Artificial Intelligence, Honolulu, HI, USA, 27 January–1 February 2019; pp. 8344–8351.
36. Gu, J.; Zhao, J.; Jiang, X.; Zhang, B.; Liu, J.; Guo, G.; Ji, R. Bayesian Optimized 1-bit Cnns. In Proceedings of the IEEE International Conference on Computer Vision, Seoul, Korea, 27 October–2 November 2019; pp. 4909–4917.