

On the Impact of Weight Quantization on Deep Neural Network Uncertainty

Shuang Liang^{1,2}, Xun Lu², Zi-Ang Liu³, Ming-Liang Wang³, Yan Lyu^{3*}, Shao-Qun Zhang^{1,2*}

¹National Key Laboratory for Novel Software Technology, Nanjing University, China

²School of Intelligent Science and Technology, Nanjing University, China

³China Mobile Zijin Innovation Institute, China

lvyanags@js.chinamobile.com, zhangsq@lamda.nju.edu.cn

Abstract

Weight quantization (WQ) is a key technique for lightweight Deep Neural Network (DNN) computations. While existing algorithms often pursue memory compression and inference acceleration with accuracy comparable to full-precision models, the effect of WQ on DNN uncertainty remains largely unexplored. In this paper, we quantify the impact of WQ on DNN uncertainty through the novel Exact Moment Propagation (EMP) uncertainty estimator. It is observed that WQ significantly increases DNN uncertainty. Based on the EMP estimator, we propose the MOMent Alignment (MOMA) to reduce WQ-induced uncertainty and preserve the accuracy of weight-quantized DNNs. Empirical results across various DNN architectures and datasets validate the effectiveness of both EMP and MOMA methods.

1 Introduction

In recent years, Weight Quantization (WQ) has attracted increasing interest in lightweight computations of Deep Neural Networks (DNNs), particularly in resource-constrained environments (Gope, Dasika, and Mattina 2019; Shen et al. 2024). Seminal studies show that WQ algorithms achieve at least 16x model compression (Li et al. 2016) and 3x training acceleration (Courbariaux, Bengio, and David 2015) while maintaining comparable accuracy to full-precision counterparts (Zhu et al. 2017; Li et al. 2022); however, the extent to which WQ techniques affect the uncertainty associated with specific outcomes remains largely unexplored.

This work focuses on quantifying the impact of WQ on DNN uncertainty. In the past decades, significant efforts have been made to quantify uncertainty in DNNs without WQ, such as Deep Ensembles (Lakshminarayanan, Pritzel, and Blundell 2017) and Moment Propagation (MP) (Wu et al. 2019). In addition, research has also examined uncertainty quantification in weight-quantized networks, such as BLRNet (Peters and Welling 2018) and BayesBiNN (Meng, Bachmann, and Khan 2020). However, a unified measure to quantify model uncertainty before and after WQ is lacking. Without this measure, one cannot fairly quantify the impact of WQ on DNN uncertainty.

In this paper, we propose the Exact Moment Propagation (EMP), a theoretical uncertainty quantification method that is independent of weight precision. Based on the EMP, we observe that quantized models exhibit greater uncertainty than their full-precision counterparts. Inspired by research on uncertainty-aware training for neural networks (Tabarisaadi et al. 2022; Einbinder et al. 2022), we propose the MOMent Alignment (MOMA), an optimization method compatible with various weight quantization training algorithms, to maintain accuracy and reduce uncertainty in quantized DNNs. This work fills the gap in quantifying and optimizing the impact of WQ on DNN uncertainty.

Our main contributions are illustrated in Figure 1(a) and further summarized as follows.

- We propose the EMP for DNN uncertainty quantification that is independent of weight precision in Algorithm 1. Empirical evidence in Subsection 5.1 shows that EMP achieves near-ideal uncertainty estimation performance, surpassing the MP, and that quantized DNNs exhibit greater uncertainty than their full-precision versions.
- We propose the MOMA for uncertainty optimization in Algorithm 2. Experiments in Subsection 5.2 show that MOMA greatly reduces the uncertainty of quantized models across different datasets.

The rest of this paper is organized as follows. Section 2 reviews related work. Section 3 and Section 4 formally introduce the EMP and MOMA methods, respectively. Section 5 presents experiments on real-world datasets to validate the effectiveness of the EMP and MOMA methods. Section 6 concludes this work.

2 Related Work

Weight Quantization. WQ aims to encode model weights into a format of k bits, where $k = 1$ and $k = 1.58$ often lead to extreme compression. For $k = 1$, the weights of DNNs are compressed into binary values, such as $\{0, 1\}$ or $\{-1, 1\}$. Seminal works include BinaryConnect (Courbariaux, Bengio, and David 2015), Binarized Neural Networks (Hubara et al. 2016), and XNOR-Net (Rastegari et al. 2016). For $k = 1.58$, the weights are compressed into ternary values $\{-1, 0, 1\}$. Early milestone works include Ternary Weight Networks (Li et al. 2016), Trained

*Corresponding authors.

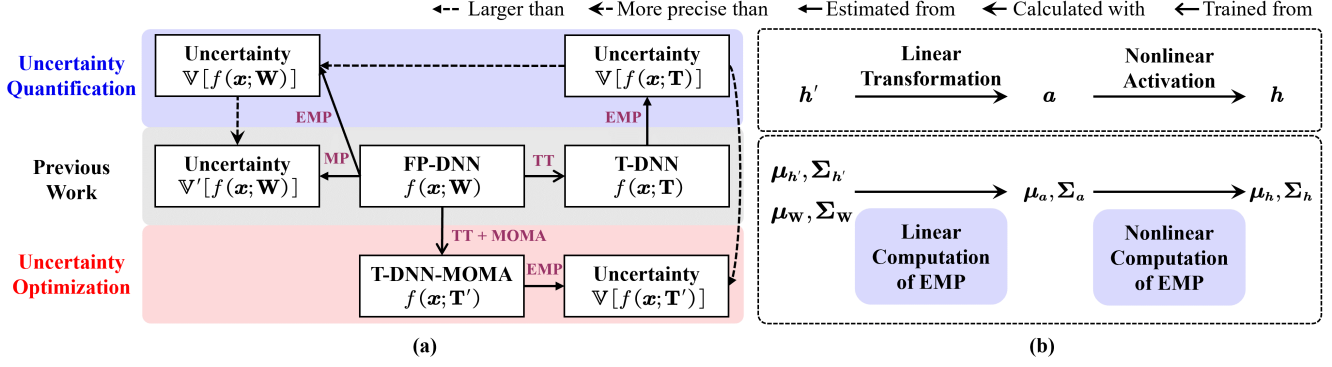


Figure 1: (a) Overview of our work, where TT is short for Ternary Training, and (b) the workflow of the EMP.

Ternary Quantization (Zhu et al. 2017), and Ternary Connect (Lin et al. 2016). In recent years, ternary WQ has been successfully applied to transformer architectures (Vaswani et al. 2017), such as TernaryBERT (Zhang et al. 2020) and Q-ViT (Li et al. 2022). TernaryCLIP (Zhang et al. 2025) achieves extremely low-cost and high-efficiency computations in image-text contrastive modeling.

Uncertainty Quantification. To quantify the uncertainty in DNNs without WQ, Deep Ensembles (Lakshminarayanan, Pritzel, and Blundell 2017) trains independent networks in parallel, estimating predictive uncertainty via negative log likelihood. PL-DNN (Bibi, Alfadly, and Ghanem 2018) and MP (Wu et al. 2019) propagate input mean and covariance through the network to derive predictive uncertainty from output covariance. For uncertainty quantification in weight-quantized DNNs, BLRNet (Peters and Welling 2018) and BayesBiNN (Meng, Bachmann, and Khan 2020) adopt the Bayesian Neural Networks paradigm (Jospin et al. 2022), sampling the posterior distribution of weights to estimate predictive uncertainty through output covariance. Despite these advances, a unified uncertainty quantification method independent of weight precision is still lacking, precluding fair comparisons of DNN uncertainty before and after WQ.

3 Uncertainty Quantification

In this section, we propose the Exact Moment Propagation (EMP) method for quantifying the uncertainty of DNNs. In contrast to previous studies, which were designed solely for full-precision models and provided either exact output covariance under strong assumptions (Bibi, Alfadly, and Ghanem 2018) or approximated output covariance (Wu et al. 2019), the proposed EMP requires only mild assumptions, improves the precision of uncertainty estimation, and extends its applicability to quantized models. The EMP method consists of the linear and nonlinear computations, whose workflow is shown in Figure 1(b).

3.1 Problem Formulation

This subsection introduces the problem formulation of WQ and uncertainty quantification of DNNs.

Ternary Weight Quantization. Throughout this work, we focus on the ternary quantization and begin with an L -layer DNN $f(\mathbf{x}; \mathbf{W})$, of which the forward propagation follows

$$\mathbf{h}^{(0)} = \mathbf{x}, \mathbf{a}^{(l)} = \mathbf{W}^{(l)} \mathbf{h}^{(l-1)} + \mathbf{b}^{(l)}, \mathbf{h}^{(l)} = \tau(\mathbf{a}^{(l)}),$$

for $l \in [L]$, where $\mathbf{W}^{(l)}$ and $\mathbf{b}^{(l)}$ are the weight matrix and bias vector of the l -th layer, respectively, τ is the nonlinear activation function, $[L]$ represents the set $\{1, 2, \dots, L\}$, and $\mathbf{W}$ denotes the set of weight matrices, i.e., $\mathbf{W} = \{\mathbf{W}^{(l)}\}_{l=1}^L$. For notational simplicity, $\mathbf{W}$ is referred to as either a weight matrix or a set, depending on the context. We define the output of the DNN as $\hat{\mathbf{y}} = \mathbf{a}^{(L)}$. In this paper, we specify τ as the ReLU activation. For notational clarity, we will omit the explicit layer index l in the following discussion and use primed symbols to represent variables from the $(l-1)$ -th layer, e.g., $\mathbf{h}' = \mathbf{h}^{(l-1)}$.

The objective of ternary quantization is to establish a ternary-weight DNN, provided a pretrained one. In the remainder of this paper, we denote $f(\mathbf{x}; \mathbf{W})$ and $f(\mathbf{x}; \mathbf{T})$ as Full-Precision DNNs (FP-DNNs) and Ternary DNNs (T-DNNs), respectively. Here, $\mathbf{T} = \{\mathbf{T}^{(l)}\}_{l=1}^L$, each element of which belongs to $\{-1, 0, 1\}$. Notice that $\mathbf{T}$ is referred to as either a weight matrix or a set, depending on the context. **Uncertainty Quantification.** Similar to Bayesian-based uncertainty quantification methods, we also treat weights as random variables and make the following assumptions.

Assumption 1. We assume that the weights within the same layer of FP-DNNs follow a Gaussian distribution, that is, $w^{(l)} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(\mu^{(l)}, (\sigma^{(l)})^2)$ for $l \in [L]$, where $w^{(l)}$ denotes an arbitrary element of matrix $\mathbf{W}^{(l)}$.

Assumption 2. We assume that the weights within the same layer of T-DNNs follow a ternary-valued distribution, that is, $t^{(l)} \stackrel{\text{i.i.d.}}{\sim} \mathcal{T}(p^{(l)}, 1 - p^{(l)} - q^{(l)}, q^{(l)})$ for $l \in [L]$, where $t^{(l)}$ denotes an arbitrary element of matrix $\mathbf{T}^{(l)}$ and $X \sim \mathcal{T}(p, 1 - p - q, q)$ denotes that the random variable X follows a discrete distribution with $P(X = -1) = p$, $P(X = 0) = 1 - p - q$, and $P(X = 1) = q$ for $p, q \in [0, 1]$.

Seminal studies (He and Fan 2019; Alawad and Ishak 2024; Zhang, Wang, and Fan 2024; Tian et al. 2025) made mild as-

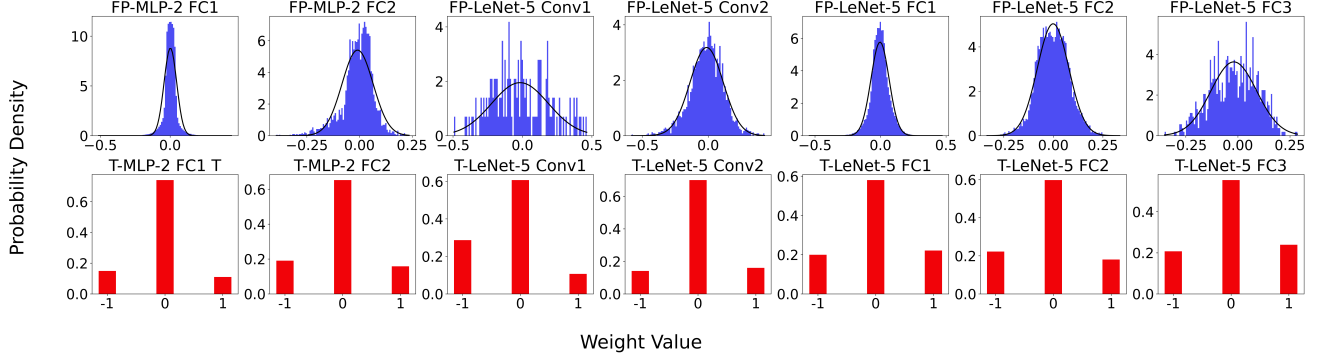


Figure 2: Empirical weight distributions of MLP-2 and LeNet-5 layer by layer.

sumptions similar to Assumption 1 for describing the weight distributions of FP-DNNs. We also conduct empirical investigations on the mildness of Assumptions 1 and 2. We first train a two-layer MLP (Rumelhart, Hinton, and Williams 1986) and a LeNet-5 (LeCun et al. 1998) with the SGD optimizer (Sutskever et al. 2013) on MNIST (LeCun et al. 1998), denoted as FP-MLP-2 and FP-LeNet-5, respectively. We then train their ternary-weight counterparts, i.e., T-MLP-2 and T-LeNet-5. Figure 2 shows the empirical weight distributions within the same layer for the concerned models, where the blue and red regions represent those of full-precision and ternary-value, respectively. It is observed that each blue region approximately follows a Gaussian distribution, hence supporting Assumption 1. In addition, it is obvious that each red region follows a ternary-valued distribution, hence supporting Assumption 2. Following this logic, we can estimate the parameters of weight distributions in Assumptions 1 and 2 from the empirical distributions.

3.2 Exact Moment Propagation

In this subsection, we introduce the EMP method. Figure 1(b) illustrates the EMP through a single layer. During feed-forward propagation, the concerned moments undergo two types of conversions. The linear transformation from $\mathbf{h}'$ to $\mathbf{a}$ induces the conversion from $(\mu_{\mathbf{h}'}, \Sigma_{\mathbf{h}'})$ to $(\mu_{\mathbf{a}}, \Sigma_{\mathbf{a}})$, which is formally expressed by the linear computation of EMP. Similarly, the nonlinear activation from $\mathbf{a}$ to $\mathbf{h}$ induces the conversion from $(\mu_{\mathbf{a}}, \Sigma_{\mathbf{a}})$ to $(\mu_{\mathbf{h}}, \Sigma_{\mathbf{h}})$, which is formally expressed by the nonlinear computation of EMP. The nonlinear computation of EMP calculates the covariance matrix $\Sigma_{\mathbf{h}}$ element by element, where each element is obtained by solving an integral. The challenge lies in solving the integrals for the non-diagonal elements, as the function being integrated becomes more complex than that in the diagonal case (Frey and Hinton 1999). Previous studies have approximated these integrals by either constraining the integration intervals (Bibi, Alfadly, and Ghanem 2018) or using asymptotic methods (Wu et al. 2019). In contrast, we solve these integrals directly by applying a recursive approach based on integration by parts, leading to an exact solution for $\Sigma_{\mathbf{h}}$. The following subsections elaborate on the completion of linear computation and nonlinear computa-

tion of the EMP.

Linear Computation of EMP. The goal of the linear computation of EMP is to provide exact expressions to calculate $(\mu_{\mathbf{a}}, \Sigma_{\mathbf{a}})$ using $(\mu_{\mathbf{h}'}, \Sigma_{\mathbf{h}'})$ and $(\mu_{\mathbf{W}}, \Sigma_{\mathbf{W}})$. We achieve this goal with the following theorem.

Theorem 3. *Let $\mathbf{h}'$ be a random vector characterized by moments $\mu_{\mathbf{h}'}$ and $\Sigma_{\mathbf{h}'}$, b_i denote a constant, and $\mathbf{W}_i$ be a random vector characterized by moments $\mu_{\mathbf{W}_i}$ and $\Sigma_{\mathbf{W}_i}$ for $i \in \{1, 2\}$. Define the scalar random variable $a_i = \mathbf{W}_i^\top \mathbf{h}' + b_i$, where $i \in \{1, 2\}$. Then, the following holds.*

$$\mathbb{E}[a_i] = \mu_{\mathbf{W}_i}^\top \mu_{\mathbf{h}'} + b_i,$$

$$\mathbb{V}(a_i) = \mu_{\mathbf{h}'}^\top \Sigma_{\mathbf{W}_i} \mu_{\mathbf{h}'} + \text{tr}(\Sigma_{\mathbf{W}_i} \Sigma_{\mathbf{h}'}) + \mu_{\mathbf{W}_i}^\top \Sigma_{\mathbf{h}'} \mu_{\mathbf{W}_i},$$

$$\begin{aligned} \text{Cov}(a_1, a_2) = & \text{tr}(\text{Cov}(\mathbf{W}_1, \mathbf{W}_2) \Sigma_{\mathbf{h}'}) + \mu_{\mathbf{W}_1}^\top \Sigma_{\mathbf{h}'} \mu_{\mathbf{W}_2} \\ & + \mu_{\mathbf{h}'}^\top \text{Cov}(\mathbf{W}_1, \mathbf{W}_2) \mu_{\mathbf{h}'} . \end{aligned}$$

Theorem 3 provides analytic expressions of the expectation and covariance of pre-activation variables. Notice that Theorem 3 is compatible with FP-DNNs and T-DNNs, leading to the following corollaries.

Corollary 4. *Under the conditions given in Theorem 3 and Assumption 1, the following holds.*

$$\mathbb{E}[a_i] = \mu \mathbf{1}^\top \mu_{\mathbf{h}'} + b_i,$$

$$\mathbb{V}(a_i) = \sigma^2 \left(\|\mu_{\mathbf{h}'}\|_2^2 + \text{tr}(\Sigma_{\mathbf{h}'}) \right) + \mu^2 \mathbf{1}^\top \Sigma_{\mathbf{h}'} \mathbf{1},$$

$$\text{Cov}(a_1, a_2) = \mu^2 \mathbf{1}^\top \Sigma_{\mathbf{h}'} \mathbf{1},$$

where $\mathbf{1}$ denotes the vector with all components equal to 1 and $\|\cdot\|_2$ denotes the L_2 norm.

Corollary 5. *Under the conditions in Theorem 3 and Assumption 2, the following holds.*

$$\mathbb{E}[a_i] = (q - p) \mathbf{1}^\top \mu_{\mathbf{h}'} + b_i,$$

$$\begin{aligned} \mathbb{V}(a_i) = & \left[(q + p) - (q - p)^2 \right] \left(\|\mu_{\mathbf{h}'}\|_2^2 + \text{tr}(\Sigma_{\mathbf{h}'}) \right) \\ & + (q - p)^2 \mathbf{1}^\top \Sigma_{\mathbf{h}'} \mathbf{1}, \end{aligned}$$

$$\text{Cov}(a_1, a_2) = (q - p)^2 \mathbf{1}^\top \Sigma_{\mathbf{h}'} \mathbf{1}.$$

Corollary 4 and Corollary 5 provide analytic expressions of the expectation and covariance of the pre-activation variable, avoiding the computation of the term $\text{Cov}(\mathbf{W}_1, \mathbf{W}_2)$

in Theorem 3, thereby enabling a more efficient linear computation of EMP than that of Theorem 3. The full proof of Theorem 3 can be found in Appendix C.

Nonlinear Computation of EMP. The goal of the nonlinear computation of EMP is to provide exact expressions to calculate (μ_h, Σ_h) using (μ_a, Σ_a) . The core is to characterize the pre-activation variable $\mathbf{a}$ according to the Law of the Unconscious Statistician (Feller 1991). In general, we make the following assumptions.

Assumption 6. Under Assumption 1, we assume that the pre-activation variables of FP-DNNs follow a Gaussian distribution, that is, $\mathbf{a}^{(l)} \sim \mathcal{N}(\mu_{\mathbf{a}^{(l)}}, \Sigma_{\mathbf{a}^{(l)}})$ for $l \in [L]$.

Assumption 7. Under Assumption 2, we assume that the pre-activation variables of T-DNNs follow a Gaussian distribution, that is, $\mathbf{a}^{(l)} \sim \mathcal{N}(\mu_{\mathbf{a}^{(l)}}, \Sigma_{\mathbf{a}^{(l)}})$ for $l \in [L]$.

We validate the mildness of Assumptions 6 and 7, and discuss the role and interdependencies of all assumptions made throughout this paper in Appendix A. Next, it suffices to analytically express (μ_h, Σ_h) under the Gaussian assumption on $\mathbf{a}$. We separately denote the Probability Density Function (PDF) and the Cumulative Distribution Function (CDF) of the standard univariate Gaussian distribution as

$$\phi(x) = \frac{1}{\sqrt{2\pi}} \exp\left(-\frac{x^2}{2}\right), \quad \Phi(x) = \int_{-\infty}^x \phi(u) du.$$

In addition, we separately denote the PDF and the CDF of the standard bivariate Gaussian distribution with correlation coefficient $\rho \in (-1, 1)$ as

$$\begin{cases} \phi_2(x_1, x_2; \rho) = \frac{1}{2\pi\sqrt{1-\rho^2}} \exp\left(\frac{x_1^2 - 2\rho x_1 x_2 + x_2^2}{-2(1-\rho^2)}\right), \\ \Phi_2(x_1, x_2; \rho) = \int_{-\infty}^{x_1} \int_{-\infty}^{x_2} \phi_2(u_1, u_2; \rho) du_1 du_2. \end{cases}$$

With the notations above, we present the following theorem.

Theorem 8. Given $(a_1, a_2)^\top \sim \mathcal{N}_2(\mu, \Sigma)$ with

$$\mu = (\mu_1, \mu_2)^\top, \quad \Sigma = \begin{pmatrix} \sigma_1^2 & \rho\sigma_1\sigma_2 \\ \rho\sigma_1\sigma_2 & \sigma_2^2 \end{pmatrix},$$

and $h_i = \tau(a_i)$ for $i \in \{1, 2\}$ where $\tau(\cdot)$ denotes the ReLU activation, then the following holds.

$$\begin{aligned} \mathbb{E}[h_i] &= \sigma_i (r_i \Phi(r_i) + \phi(r_i)), \\ \mathbb{V}[h_i] &= \sigma_i^2 [(1 + r_i^2) \Phi(r_i) + r_i \phi(r_i)] - \mathbb{E}[h_i]^2, \\ \text{Cov}(h_1, h_2) &= \sigma_1 \sigma_2 [r_2 \phi(r_1) \Phi(c_2) + r_1 \phi(r_2) \Phi(c_1) \\ &\quad + (r_1 r_2 + \rho) \Phi_2(r_1, r_2; \rho) + \sqrt{1 - \rho^2} \phi(r_2) \phi(c_1)] \\ &\quad - \mathbb{E}[Z_1] \mathbb{E}[Z_2], \end{aligned}$$

in which $r_i = \mu_i/\sigma_i$, $c_i = (r_i - \rho r_{3-i})/\sqrt{1 - \rho^2}$ for $i \in \{1, 2\}$, $\phi(\cdot)$ and $\Phi(\cdot)$ denote the PDF and the CDF of the standard univariate normal distribution, respectively, and $\Phi_2(\cdot, \cdot; \rho)$ denotes the CDF of standard bivariate normal distribution with correlation coefficient ρ .

Theorem 8 provides analytic expressions of expectation and covariance of $\tau(\mathbf{a})$, completing the nonlinear computation

Algorithm 1: EMP for T-DNNs

Input: Sample distribution μ_x, Σ_x , weight distribution parameters $\{p^{(l)}, q^{(l)}\}_{l=1}^L$

Output: Output covariance $\Sigma_{\hat{y}}$

Procedure:

- 1: Set $\mu_h^{(0)} = \mu_x, \Sigma_h^{(0)} = \Sigma_x$
 - 2: **for** $1 \leq l \leq L$ **do**
 - 3: $(\mu_a^{(l)}, \Sigma_a^{(l)}) \leftarrow (\mu_h^{(l-1)}, \Sigma_h^{(l-1)})$ and $(p^{(l)}, q^{(l)})$ by Corollary 5
 - 4: $(\mu_h^{(l)}, \Sigma_h^{(l)}) \leftarrow (\mu_a^{(l)}, \Sigma_a^{(l)})$ by Theorem 8
 - 5: **end for**
 - 6: $\Sigma_{\hat{y}} = \Sigma_a^{(L)}$
-

of EMP. Notice that Theorem 8 is straightforwardly compatible with both FP-DNNs and T-DNNs. To the best of our knowledge, Theorem 8 provides the first exact solution to the covariance Σ_h of a Gaussian vector passed through the ReLU activation under mild assumptions; in contrast, the previous methods provided either exact Σ_h under strong assumptions (Bibi, Alfadly, and Ghanem 2018) or approximated Σ_h with $\mathbf{h} = \tau(\mathbf{a})$ (Wu et al. 2019). The full proof of Theorem 8 can be found in Appendix D.

To summarize the computations above, we can complete the EMP method for uncertainty quantification. Algorithm 1 lists the procedure of the EMP method for quantifying the uncertainty of T-DNNs, where Step 3 and Step 4 represent the linear computation and the nonlinear computation of the EMP, respectively. Algorithm 1 also adapts to FP-DNNs with two modifications: replacing the input $\{p^{(l)}, q^{(l)}\}_{l=1}^L$ with $\{\mu^{(l)}, \sigma^{(l)}\}_{l=1}^L$ and substituting Corollary 5 in Step 3 with Corollary 4.

4 Uncertainty Optimization

In this section, we propose an optimization method termed MOMent Alignment (MOMA) for improving the accuracy and reducing the uncertainty of T-DNNs. The MOMA is compatible with various quantization training algorithms for DNNs. Here, we also adopt ternary quantization as an example and draw the key ideas from Li et al. (2016) and Alemдар et al. (2017) to achieve ternary training. We detail the optimization procedure in Subsection 4.1, and then formally propose the corresponding MOMA in Subsection 4.2.

4.1 Ternary Training of T-DNNs

Following the supervised learning paradigm, the optimization problem of training a T-DNN f can be formulated as

$$\min_{\mathbf{T}} \mathcal{L}(\mathbf{y}, f(\mathbf{x}; \mathbf{T})), \quad (1)$$

where $(\mathbf{x}, \mathbf{y})$ denotes the pair of input and target variables, and $\mathcal{L}(\cdot, \cdot)$ denotes the loss function. Let $s \in \mathbb{N}$ indicate the training iteration. The weight-updating procedure is listed as

$$\begin{cases} \mathbf{T}_{s+1} = \text{sgn}_{\gamma_s}(\mathbf{W}_{s+1}), \\ \mathbf{W}_{s+1} = \mathbf{W}_s + \eta \frac{\partial \mathcal{L}(\mathbf{y}, f(\mathbf{x}; \mathbf{T}_s))}{\partial \mathbf{W}_s}, \end{cases} \quad (2)$$

where η denotes the learning rate and $\text{sgn}_{\gamma_s}(\cdot)$ is the partition function. Here, we employ the partition function used in Li et al. (2016), that is,

$$t = \text{sgn}_{\gamma_s}(w) = \begin{cases} 1, & w > \gamma_s, \\ 0, & -\gamma_s \leq w \leq \gamma_s, \\ -1, & w < -\gamma_s, \end{cases} \quad (3)$$

where t and w denote an arbitrary element of the weight matrices $\mathbf{T}_s$ and $\mathbf{W}_s$, respectively, and $\gamma_s \in \mathbb{R}^+$ is a partition factor defined as

$$\gamma_s = \frac{1}{N} \|\mathbf{W}_s\|_1. \quad (4)$$

The gradient-related term $\partial \mathcal{L}(\mathbf{y}, f(\mathbf{x}; \mathbf{T}_s)) / \partial \mathbf{W}_s$ in Eq. (2) is computed via surrogate gradients (Bengio, Léonard, and Courville 2013) as

$$\begin{cases} \frac{\partial \mathcal{L}(\mathbf{y}, f(\mathbf{x}; \mathbf{T}_s))}{\partial \mathbf{W}_s} = \frac{\partial \mathcal{L}(\mathbf{y}, f(\mathbf{x}; \mathbf{T}_s))}{\partial \mathbf{T}_s} \frac{\partial \mathbf{T}_s}{\partial \mathbf{W}_s}, \\ \frac{\partial \mathbf{T}_s}{\partial \mathbf{W}_s} = \frac{\partial \text{sgn}_{\gamma_s}(\mathbf{W}_s)}{\partial \mathbf{W}_s} = 1. \end{cases} \quad (5)$$

In summary, one can obtain a T-DNN by initializing $\mathbf{W}_0$ with the pretrained full-precision weights, and following the optimization procedure listed in Eq. (2), Eq. (4), and Eq. (5).

4.2 Moment Alignment

This subsection proposes the MOMA optimization for improving the accuracy and reducing the uncertainty of T-DNNs, ternarized from a pretrained FP-DNN. The key idea of MOMA is to formulate an optimization problem by aligning the first and second central moments of T-DNN weights to zero. The following theorem supports our line of thought.

Theorem 9. *Under Assumption 2 and Assumption 7, if $q - p = 0$ and $p + q - (q - p)^2 = 0$ hold for the ternary-valued distribution parameters (p, q) of each layer, then the output covariance satisfies $\Sigma_{\hat{\mathbf{y}}} = \mathbf{0}$.*

Theorem 9 demonstrates that there exists an optimal solution, that is, the uncertainty of T-DNNs can be reduced to zero if one aligns the concerned moments to zero, which provides a theoretical guarantee for MOMA. The full proof of Theorem 9 can be found in Appendix E. Inspired by Theorem 9, we formulate the MOMA optimization as

$$\begin{cases} \tilde{\gamma}_s = \arg \min_{\gamma} \{ |g_3(\gamma, \mathbf{W}_s)| + \lambda |g_4(\gamma, \mathbf{W}_s)| \}, \\ \text{s.t. } \gamma \in ((1 - c)\gamma_s, (1 + c)\gamma_s). \end{cases} \quad (6)$$

Algorithm 2: Ternary Training with MOMA for T-DNNs

Input: Pretrained full-precision weights $\mathbf{W}_0$, the maximum number of training epochs S_0

Output: Ternary weight $\mathbf{T}$

Procedure:

- 1: **for** $0 \leq s \leq S_0$ **do**
 - 2: $\mathbf{W}_{s+1} = \mathbf{W}_s + \eta \partial \mathcal{L}(\mathbf{y}, f(\mathbf{x}; \mathbf{T}_s)) / \partial \mathbf{T}_s$ in Eq. (2)
 - 3: $\gamma_s = \frac{1}{N} \|\mathbf{W}_s\|_1$ in Eq. (4)
 - 4: $\tilde{\gamma}_s \leftarrow$ solving Eq. (6)
 - 5: $\mathbf{T}_{s+1} = \text{sgn}_{\tilde{\gamma}_s}(\mathbf{W}_{s+1})$ in Eq. (3).
 - 6: **end for**
-

where $c \in (0, 1)$, $\lambda \in \mathbb{R}^+$, g_3 is an estimator for $q - p$, and g_4 is an estimator for $p + q - (q - p)^2$. The search interval is designed to be a neighborhood of the conventional partition factor of Eq. (4) to maintain accuracy, while the objective function is designed to align the concerned moments to zero for uncertainty reduction. The estimators in Eq. (6) are defined as follows.

$$\begin{cases} g_1(\gamma, \mathbf{W}_s) = \sum_{i,j} \chi_{(-\infty, -\gamma)}(\mathbf{W}_{ij}) / N, \\ g_2(\gamma, \mathbf{W}_s) = \sum_{i,j} \chi_{(\gamma, +\infty)}(\mathbf{W}_{ij}) / N, \\ g_3(\gamma, \mathbf{W}_s) = g_2(\gamma, \mathbf{W}_s) - g_1(\gamma, \mathbf{W}_s), \\ g_4(\gamma, \mathbf{W}_s) = g_1(\gamma, \mathbf{W}_s) + g_2(\gamma, \mathbf{W}_s) - (g_3(\gamma, \mathbf{W}_s))^2, \end{cases}$$

where g_1 is an estimator for p , g_2 is an estimator for q , $\mathbf{W}_{ij}$ denotes the (i, j) -th element of $\mathbf{W}$, and $\chi_A(\cdot)$ denotes the indicator function of a concerned interval A that satisfies $\chi_A(x) = 1$ if $x \in A$ and $\chi_A(x) = 0$ if $x \notin A$.

Algorithm 2 details the calculation procedure of the MOMA optimization for T-DNN training. In contrast to conventional T-DNN training algorithms, MOMA involves an additional search for the partition factor, as detailed in Step 5. Hence, our proposed MOMA method can be incorporated into other quantization training algorithms that use a threshold-based partition function like Eq. (3).

5 Experiments

In this section, we conduct experiments to verify the effectiveness of the proposed EMP and MOMA methods. The experiments are performed to answer the following questions.

- Q1. Is the proposed EMP method comparable to the classical uncertainty estimation methods?
- Q2. To what extent, if at all, does weight ternarization increase the uncertainty of DNNs?
- Q3. Whether and to what extent does the proposed MOMA method reduce the uncertainty of T-DNNs?

All experiments were conducted on MNIST (LeCun et al. 1998), CIFAR-10, and CIFAR-100 (Krizhevsky 2009) datasets. The concerned architectures include MLP-2¹ (Rumelhart, Hinton, and Williams 1986), LeNet-5 (LeCun et al. 1998), VGG-13 (Simonyan 2014), and ResNet-18 (He et al. 2016). We denote the full-precision DNNs as FP-DNNs, such as FP-MLP-2. The ternary-weight DNNs trained using the procedures outlined in Subsection 4.1 are referred to as T-DNNs, such as T-MLP-2. The ternary-weight DNNs trained with Algorithm 2 are denoted as T-DNNs-MOMA, such as T-MLP-2-MOMA. The proposed methods and contenders are implemented in PyTorch and SciPy. All experiments were conducted on a 13th gen Intel(R) Core(TM) i9-13900KF CPU and an NVIDIA RTX 4090 GPU. Experiments were repeated 5 times with different initialization seeds, reporting the mean α and standard deviation β as α_β .

¹MLP-2 denotes a two-layer fully-connected architecture.

Datasets	o	FP-LeNet-5			T-LeNet-5		
		EMP	MP	PL-DNN	EMP	MP	PL-DNN
MNIST	0	1.0007 ± 0.0101	1.1809 ± 0.0217	1.2772 ± 0.0153	1.0023 ± 0.0109	—	—
	1	0.9999 ± 0.0103	1.1793 ± 0.0238	1.2759 ± 0.0139	1.0023 ± 0.0121	—	—
	2	0.9990 ± 0.0118	1.1794 ± 0.0239	1.2748 ± 0.0170	1.0021 ± 0.0117	—	—
	3	1.0008 ± 0.0104	1.1808 ± 0.0215	1.2770 ± 0.0153	1.0022 ± 0.0108	—	—
	4	1.0011 ± 0.0096	1.1807 ± 0.0239	1.2775 ± 0.0136	1.0023 ± 0.0107	—	—
	5	1.0010 ± 0.0106	1.1815 ± 0.0237	1.2773 ± 0.0149	1.0022 ± 0.0109	—	—
	6	1.0011 ± 0.0096	1.1818 ± 0.0246	1.2774 ± 0.0130	1.0023 ± 0.0105	—	—
	7	0.9996 ± 0.0099	1.1803 ± 0.0243	1.2755 ± 0.0143	1.0021 ± 0.0101	—	—
	8	1.0014 ± 0.0098	1.1814 ± 0.0218	1.2422 ± 0.0140	1.0021 ± 0.0107	—	—
	9	0.9999 ± 0.0100	1.1801 ± 0.0225	1.2759 ± 0.0142	1.0019 ± 0.0106	—	—
CIFAR-10	0	1.0098 ± 0.0099	1.5029 ± 0.0177	1.6167 ± 0.0334	1.0141 ± 0.0113	—	—
	1	1.0092 ± 0.0105	1.5020 ± 0.0207	1.6156 ± 0.0299	1.0136 ± 0.0115	—	—
	2	1.0090 ± 0.0127	1.5017 ± 0.0218	1.6154 ± 0.0358	1.0141 ± 0.0110	—	—
	3	1.0100 ± 0.0110	1.5031 ± 0.0189	1.6169 ± 0.0338	1.0141 ± 0.0115	—	—
	4	1.0106 ± 0.0098	1.5041 ± 0.0179	1.6179 ± 0.0307	1.0137 ± 0.0116	—	—
	5	1.0099 ± 0.0109	1.5030 ± 0.0200	1.6167 ± 0.0311	1.0139 ± 0.0118	—	—
	6	1.0103 ± 0.0094	1.5036 ± 0.0185	1.6174 ± 0.0282	1.0139 ± 0.0109	—	—
	7	1.0090 ± 0.0099	1.5017 ± 0.0184	1.6153 ± 0.0307	1.0139 ± 0.0120	—	—
	8	1.0103 ± 0.0093	1.5036 ± 0.0179	1.6173 ± 0.0312	1.0140 ± 0.0118	—	—
	9	1.0095 ± 0.0100	1.5025 ± 0.0179	1.6162 ± 0.0314	1.0141 ± 0.0109	—	—

Table 1: Performance of the EMP and contenders for LeNet-5 on MNIST and CIFAR-10. Results for other architectures and datasets are detailed in Appendix F.

5.1 Verifications on EMP

This subsection demonstrates the effectiveness of the proposed EMP method. Let $\mathbf{x} \in D$ denote a test sample from the test set D . In Algorithm 1, we set $\boldsymbol{\mu}_{\mathbf{x}} = \mathbf{x}$ and $\boldsymbol{\Sigma}_{\mathbf{x}} = \mathbf{0}$.

To compare the performance of various uncertainty estimators, we follow the evaluation metric in Bibi, Alfadly, and Ghanem (2018), where they treat the classical Monte Carlo (MC) estimator (Jospin et al. 2022) as the ground truth, with other estimators being evaluated based on their closeness to it. Formally, we define the variance ratio as $r_o(\mathbf{x}) = \mathbb{V}_o^{\text{MC}}(\mathbf{x}) / \mathbb{V}_o^{\text{E}}(\mathbf{x})$, where $\mathbb{V}_o^{\text{MC}}(\mathbf{x})$ and $\mathbb{V}_o^{\text{E}}(\mathbf{x})$ denote the o -th element of the output variance vector obtained by the MC method and by an estimator, respectively, for $o \in \{0, 1, \dots, 9\}$. A ratio $r_o(\mathbf{x})$ closer to 1 indicates higher precision of an estimator on $\mathbf{x}$. We then compute the mean and standard deviation of $r_o(\mathbf{x})$ over D as

$$\mathbb{E}[r_o(\mathbf{x})] = \sum_{\mathbf{x} \in D} r_o(\mathbf{x}) / |D| ,$$

$$\sqrt{\mathbb{V}[r_o(\mathbf{x})]} = \sqrt{\sum_{\mathbf{x} \in D} (r_o(\mathbf{x}) - \mathbb{E}[r_o(\mathbf{x})])^2 / |D|} .$$

Then we use $\mathbb{E}[r_o(\mathbf{x})] \pm \sqrt{\mathbb{V}[r_o(\mathbf{x})]}$ as a metric, which is the same as that in Bibi, Alfadly, and Ghanem (2018), to measure the uncertainty estimation performance of an estimator. A mean $\mathbb{E}[r_o(\mathbf{x})]$ closer to 1 indicates higher precision on average, while a standard deviation $\sqrt{\mathbb{V}[r_o(\mathbf{x})]}$ closer to 0 indicates greater robustness on average.

Table 1 shows the uncertainty estimation performance of EMP and its contenders, including PL-DNN (Bibi, Alfadly,

and Ghanem 2018) and MP (Wu et al. 2019) for LeNet-5 on MNIST and CIFAR-10. It is observed that the EMP method achieves near-ideal performance and significantly outperforms its contenders. Experimental results for other architectures and datasets, detailed in Appendix F, further support this observation. These findings demonstrate that the proposed EMP estimator is comparable to the MC estimator and more precise than its contenders, thus answering Q1.

5.2 Verifications on MOMA

This experiment demonstrates the effectiveness of the MOMA method. We first analyze the uncertainty difference between DNNs with and without weight ternarization to take a closer look at Q2. We begin this analysis by defining two variance ratios as

$$\tilde{r}_o(\mathbf{x}) = \mathbb{V}_o^{\text{E}}(\mathbf{x}; \text{T-DNN}) / \mathbb{V}_o^{\text{E}}(\mathbf{x}; \text{FP-DNN}) ,$$

$$\tilde{r}'_o(\mathbf{x}) = \mathbb{V}_o^{\text{E}}(\mathbf{x}; \text{T-DNN-MOMA}) / \mathbb{V}_o^{\text{E}}(\mathbf{x}; \text{FP-DNN}) ,$$

where $\mathbb{V}_o^{\text{E}}(\mathbf{x}; \text{Model})$ represents the variance, i.e., $\mathbb{V}_o^{\text{E}}(\mathbf{x})$, computed using the corresponding model. A ratio greater than 1 indicates that weight ternarization increases DNN uncertainty by the corresponding ratio.

Figure 3 shows the concerned uncertainty difference on MLP-2 by providing the histogram of the two ratios. Due to space limitations, we only plot the results for $o \in \{0, 1, \dots, 5\}$ and the ratio range of $[0, 3000]$, containing 99% of the empirical points for the variance ratios. Full results are given in Appendix G. The blue and red regions represent the histograms of $\tilde{r}_o(\mathbf{x})$ and $\tilde{r}'_o(\mathbf{x})$, respectively. It is observed that $\tilde{r}_o(\mathbf{x}) > 1.4$ and $\tilde{r}'_o(\mathbf{x}) > 1.4$ for all

Models	Accuracy	Uncertainty	Reduction
FP-MLP-2	0.9761 _{0.0002}	$1.4697_{0.0082} \times 10^3$	68.08 _{0.1389} %
T-MLP-2	0.9752 _{0.0005}	$1.2572_{0.0236} \times 10^5$	
T-MLP-2-MOMA	0.9751 _{0.0006}	$4.0129_{0.0039} \times 10^4$	
FP-LeNet-5	0.9911 _{0.0001}	$1.7728_{0.0232} \times 10^3$	98.55 _{0.0025} %
T-LeNet-5	0.9903 _{0.0003}	$2.8113_{0.0490} \times 10^5$	
T-LeNet-5-MOMA	0.9885 _{0.0006}	$4.0906_{0.1433} \times 10^3$	
FP-VGG-13	0.9985 _{0.0001}	$4.6029_{0.2302} \times 10^{-1}$	20.86 _{2.5366} %
T-VGG-13	0.9954 _{0.0005}	$1.4500_{0.0210} \times 10^1$	
T-VGG-13-MOMA	0.9957 _{0.0005}	$1.1463_{0.0198} \times 10^1$	
FP-ResNet-18	0.9963 _{0.0001}	$6.9757_{0.3743} \times 10^{-3}$	44.27 _{4.4212} %
T-ResNet-18	0.9938 _{0.0002}	$4.2808_{0.1772} \times 10^{-1}$	
T-ResNet-18-MOMA	0.9951 _{0.0002}	$2.4067_{0.0501} \times 10^{-1}$	

Table 2: Accuracy and uncertainty of concerned models on MNIST. Results for other datasets are detailed in Appendix G.

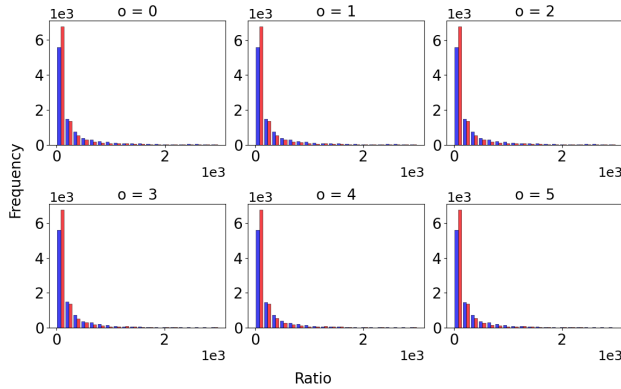


Figure 3: Histograms of the concerned variance ratios.

$\mathbf{x} \in D$, indicating that DNN uncertainty increases after weight ternarization, regardless of whether MOMA is used. This observation answers Q2. Furthermore, the red bars are taller at lower x-values and shorter at higher x-values, indicating that the proposed MOMA method effectively reduces the concerned variance ratio.

To further measure the uncertainty reduction performance of the proposed MOMA method, we define the uncertainty of a model on a dataset D as

$$\text{Uncertainty} = \frac{1}{|D|} \sum_{\mathbf{x} \in D} \text{tr}(\Sigma_{\hat{\mathbf{y}}(\mathbf{x})}).$$

Table 2 shows the uncertainty reduction performance of MOMA across various DNN architectures on MNIST by comparing the uncertainty of the T-DNNs and T-DNNs-MOMA. It is observed that the MOMA greatly reduces the uncertainty of T-DNNs on MNIST. Experimental results for CIFAR-10 and CIFAR-100, provided in Appendix G, further support this observation. This finding demonstrates that the proposed MOMA can effectively reduce the uncertainty of T-DNNs, thus answering Q3.

6 Conclusions

In this paper, we proposed the EMP, a theoretical uncertainty quantification method that works independently of weight precision. Experimental results showed that T-DNNs exhibit significantly greater uncertainty than FP-DNNs across varying ratios. Inspired by the EMP estimator, we proposed MOMA, an optimization method for reducing the uncertainty of T-DNNs. Empirical results validated the effectiveness of both EMP and MOMA.

Acknowledgements

This research was supported by the Jiangsu Provincial Natural Science Foundation Youth Project (BK20230782) and Nanjing University-China Mobile Communications Group Co., Ltd. Joint Institute.

References

- Alawad, M.; and Ishak, M. 2024. Probabilistic Bayesian Neural Networks for Efficient Inference. In *Proceedings of the 34th Great Lakes Symposium on VLSI*, 724–729.
- Alemdar, H.; Leroy, V.; Prost-Boucle, A.; and Pétrot, F. 2017. Ternary neural networks for resource-efficient AI applications. In *Proceedings of the 30th International Joint Conference on Neural Networks*, 2547–2554.
- Bengio, Y.; Léonard, N.; and Courville, A. 2013. Estimating or propagating gradients through stochastic neurons for conditional computation. *arXiv preprint arXiv:1308.3432*.
- Bibi, A.; Alfadly, M.; and Ghanem, B. 2018. Analytic expressions for probabilistic moments of PL-DNN with Gaussian input. In *Proceedings of the 31st IEEE Conference on Computer Vision and Pattern Recognition*, 9099–9107.
- Courbariaux, M.; Bengio, Y.; and David, J.-P. 2015. Binaryconnect: Training deep neural networks with binary weights during propagations. In *Advances in Neural Information Processing Systems* 28, 3123–3131.
- Einbinder, B.-S.; Romano, Y.; Sesia, M.; and Zhou, Y. 2022. Training uncertainty-aware classifiers with conformalized

- deep learning. In *Advances in Neural Information Processing Systems* 35, 22380–22395.
- Feller, W. 1991. *An introduction to probability theory and its applications*. John Wiley & Sons.
- Frey, B.; and Hinton, G. 1999. Variational learning in nonlinear Gaussian belief networks. *Neural Computation*, 11(1): 193–213.
- Gope, D.; Dasika, G.; and Mattina, M. 2019. Ternary hybrid neural-tree networks for highly constrained iot applications. *Proceedings of Machine Learning and Systems*, 1: 190–200.
- He, K.; Zhang, X.; Ren, S.; and Sun, J. 2016. Deep residual learning for image recognition. In *Proceedings of the 29th IEEE Conference on Computer Vision and Pattern Recognition*, 770–778.
- He, Z.; and Fan, D. 2019. Simultaneously optimizing weight and quantizer of ternary neural network using truncated gaussian approximation. In *Proceedings of the 32nd IEEE Conference on Computer Vision and Pattern Recognition*, 11438–11446.
- Hubara, I.; Courbariaux, M.; Soudry, D.; El-Yaniv, R.; and Bengio, Y. 2016. Binarized neural networks. *Advances in Neural Information Processing Systems* 29, 4107–4115.
- Jospin, L.-V.; Laga, H.; Boussaid, F.; Buntine, W.; and Benamoun, M. 2022. Hands-on Bayesian neural networks—A tutorial for deep learning users. *IEEE Computational Intelligence Magazine*, 17(2): 29–48.
- Krizhevsky, A. 2009. Learning multiple layers of features from tiny images. *Master’s thesis*.
- Lakshminarayanan, B.; Pritzel, A.; and Blundell, C. 2017. Simple and scalable predictive uncertainty estimation using deep ensembles. In *Proceedings of the 31st Advances in Neural Information Processing Systems* 30, 6405–6416.
- LeCun, Y.; Bottou, L.; Bengio, Y.; and Haffner, P. 1998. Gradient-based learning applied to document recognition. *Proceedings of the IEEE*, 86(11): 2278–2324.
- Li, F.; Liu, B.; Wang, X.; Zhang, B.; and Yan, J. 2016. Ternary weight networks. *arXiv preprint arXiv:1605.04711*.
- Li, Y.; Xu, S.; Zhang, B.; Cao, X.; Gao, P.; and Guo, G. 2022. Q-vit: Accurate and fully quantized low-bit vision transformer. In *Advances in neural information processing systems* 35, 34451–34463.
- Lin, Z.; Courbariaux, M.; Memisevic, R.; and Bengio, Y. 2016. Neural networks with few multiplications. In *Proceedings of the 4th International Conference on Learning Representations*.
- Meng, X.; Bachmann, R.; and Khan, M.-E. 2020. Training binary neural networks using the Bayesian learning rule. In *Proceedings of the 37th International Conference on Machine Learning*, 6852–6861.
- Peters, J.; and Welling, M. 2018. Probabilistic binary neural networks. *arXiv preprint arXiv:1809.03368*.
- Rastegari, M.; Ordonez, V.; Redmon, J.; and Farhadi, A. 2016. Xnor-net: Imagenet classification using binary convolutional neural networks. In *Proceedings of the 14th European Conference on Computer Vision*, 525–542.
- Rumelhart, D.; Hinton, G.; and Williams, R. 1986. Learning representations by back-propagating errors. *Nature*, 323(6088): 533–536.
- Shen, X.; Dong, P.; Lu, L.; Kong, Z.; Li, Z.; Lin, M.; Wu, C.; and Wang, Y. 2024. Agile-Quant: Activation-Guided Quantization for Faster Inference of LLMs on the Edge. In *Proceedings of the 38th AAAI Conference on Artificial Intelligence*, 18944–18951.
- Simonyan, K. 2014. Very deep convolutional networks for large-scale image recognition. *arXiv preprint arXiv:1409.1556*.
- Sutskever, I.; Martens, J.; Dahl, G.; and Hinton, G. 2013. On the importance of initialization and momentum in deep learning. In *Proceedings of the 30th International Conference on Machine Learning*, 1139–1147.
- Tabarisaadi, P.; Khosravi, A.; Nahavandi, S.; Shafie-Khah, M.; and Catalão, J. 2022. An optimized uncertainty-aware training framework for neural networks. *IEEE Transactions on Neural Networks and Learning Systems*, 35(5): 6928–6935.
- Tian, Y.-M.; Liang, S.; Zhang, S.-Q.; and Fan, F.-L. 2025. Depth-induced NTK: Bridging Over-parameterized Neural Networks and Deep Neural Kernels. *arXiv preprint arXiv:2511.05585*.
- Vaswani, A.; Shazeer, N.; Parmar, N.; Uszkoreit, J.; Jones, L.; Gomez, A.; Kaiser, Ł.; and Polosukhin, I. 2017. Attention is all you need. In *Advances in neural information processing systems* 30, 5998–6008.
- Wu, A.; Nowozin, S.; Meeds, E.; Turner, R.; Hernandez-Lobato, J.; and Gaunt, A. 2019. Deterministic variational inference for robust Bayesian neural networks. In *Proceedings of the 7th International Conference on Learning Representations*.
- Zhang, S.-H.; Tang, W.-C.; Wu, C.; Hu, P.; Li, N.; Zhang, L.-J.; Zhang, Q.; and Zhang, S.-Q. 2025. TernaryCLIP: Efficiently Compressing Vision-Language Models with Ternary Weights and Distilled Knowledge. *arXiv:2510.21879*.
- Zhang, S.-Q.; Wang, F.; and Fan, F.-L. 2024. Neural Network Gaussian Processes by Increasing Depth. *IEEE Transactions on Neural Networks and Learning Systems*, 35(2): 2881–2886.
- Zhang, W.; Hou, L.; Yin, Y.; Shang, L.; Chen, X.; Jiang, X.; and Liu, Q. 2020. Ternarybert: Distillation-aware ultra-low bit bert. In *Proceedings of the 25th Empirical Methods in Natural Language Processing*, 509–521.
- Zhu, C.; Han, S.; Mao, H.; and Dally, W. 2017. Trained ternary quantization. In *Proceedings of the 5th International Conference on Learning Representations*, 1–10.

Technical Appendix

This Technical Appendix provides the supplementary materials for our work “On the Impact of Weight Quantization on Deep Neural Network Uncertainty”, constructed according to the corresponding sections therein.

A About the Assumptions

This appendix validates Assumptions 6 and 7, and discusses the role and interdependencies of all assumptions made throughout this paper.

Seminal studies (Bibi, Alfadly, and Ghanem 2018; Wu et al. 2019) made similar assumptions to Assumption 6 for uncertainty quantification. To further verify the mildness of Assumption 6 and Assumption 7, we conduct the following experiment. We sample weights according to Assumption 1 or Assumption 2, then propagate the networks 10,000 times on a fixed MNIST instance to observe the empirical distributions of hidden layer and pre-activation variable for an MLP-2.

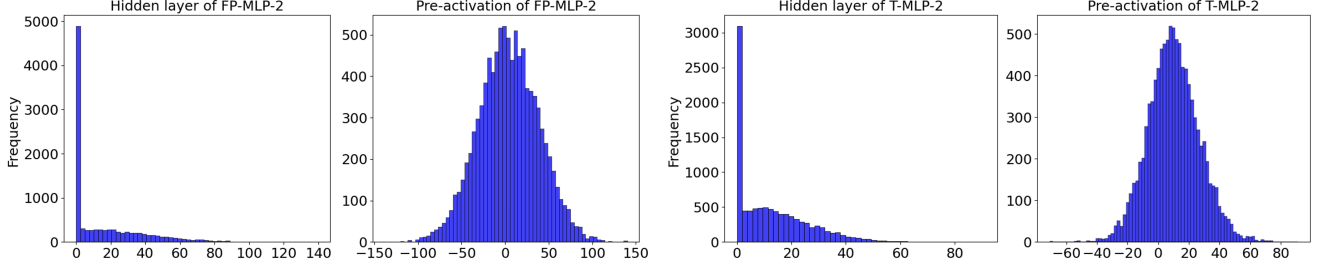


Figure 4: Empirical distributions of the hidden layer and pre-activation variables for FP-MLP-2 and T-MLP-2.

Figure 4 displays the empirical distributions of the pre-activation variable a and the hidden layer h , where the weight vector $\mathbf{W}$ satisfies Assumption 1 for FP-MLP-2 and Assumption 2 for T-MLP-2. In both cases, the empirical distribution of a closely approximates a Gaussian distribution even though h does not, indicating that Assumption 6 and Assumption 7 empirically follow from Assumption 1 and Assumption 2, respectively.

Discussion. The assumptions of weights are listed in Assumptions 1 and 2, where the weights within the same layer are i.i.d. following a Gaussian distribution for FP-DNNs and a ternary-valued distribution for T-DNNs. These assumptions are key to deriving Corollary 4, Corollary 5, and Theorem 9, providing theoretical guarantees for the effectiveness of EMP and MOMA, such as closed-form solutions for linear EMP computations and moment alignment. Besides, Assumptions 6 and 7 empirically follow from Assumptions 1 and 2. Assumptions 6 and 7 are key to deriving Theorem 8, providing theoretical guarantees for the effectiveness of EMP, such as closed-form solutions for nonlinear EMP computations.

Notice that EMP and MOMA are implemented by Algorithm 1 and Algorithm 2, respectively, which make no assumptions about the input distribution in practice. We set the input distribution as a single-point distribution in our experiments, that is, $\mu_x = x$ and $\Sigma_x = \mathbf{0}$ for an input x .

B About the limitation and complexity

B.1 EMP

About the complexity and acceleration. Uncertainty quantification has suffered from high computational demands. We address this by proposing two weight distribution assumptions: a Gaussian distribution for weights in the same layer of FP-DNNs in Assumption 1 and a ternary-valued distribution for T-DNNs in Assumption 2. These assumptions reduce variance-related uncertainty quantification complexity. While traditional BNNs require storage of $2m^2L$ distribution parameters for an MLP with depth L and width m , EMP requires only $2L$ parameters, substantially enhancing efficiency through this significant parameter reduction.

About the limitation. The EMP estimator is specifically designed for DNNs with ReLU activations.

B.2 MOMA

About the complexity. The key idea of MOMA is transforming Eq. (4) into the optimization problem Eq. (6). Therefore, the additional computational complexity introduced by MOMA comes solely from solving the optimization problem Eq. (6). This optimization problem can be efficiently solved using zero-order optimization algorithms such as grid search.

About the limitation. The MOMA is specifically designed for threshold-based quantized training algorithms, where quantized weights are obtained through the partition function Eq. (3) or its variants.

C Full Proof of Theorem 3

The key idea to calculate $\mathbb{E}[a_i]$ is to use the law of total expectation as

$$\begin{aligned}
 \mathbb{E}[a_i] &= \mathbb{E}_{\mathbf{h}'} [\mathbb{E}_{\mathbf{W}} [a_i | \mathbf{h}']] \\
 &= \mathbb{E}_{\mathbf{h}'} [\mathbb{E}_{\mathbf{W}} [\mathbf{W}_i^\top \mathbf{h}' + b_i | \mathbf{h}']] \\
 &= \mathbb{E}_{\mathbf{h}'} [\mathbb{E}_{\mathbf{W}} [\mathbf{W}_i^\top] \mathbf{h}' + b_i] \\
 &= \mathbb{E}_{\mathbf{h}'} [\boldsymbol{\mu}_{\mathbf{W}_i}^\top \mathbf{h}' + b_i] \\
 &= \boldsymbol{\mu}_{\mathbf{W}_i}^\top \boldsymbol{\mu}_{\mathbf{h}'} + b_i .
 \end{aligned} \tag{7}$$

The key idea to calculate $\mathbb{V}[a_i]$ is to utilize the law of total variance as

$$\begin{aligned}
 \mathbb{V}(a_i) &= \mathbb{E}_{\mathbf{h}'} [\mathbb{V}_{\mathbf{W}} (a_i | \mathbf{h}')] + \mathbb{V}_{\mathbf{h}'} (\mathbb{E}_{\mathbf{W}} [a_i | \mathbf{h}']) \\
 &= \mathbb{E}_{\mathbf{h}'} [\mathbb{V}_{\mathbf{W}} (\mathbf{W}_i^\top \mathbf{h}' + b_i | \mathbf{h}')] + \mathbb{V}_{\mathbf{h}'} (\mathbb{E}_{\mathbf{W}} [\mathbf{W}_i^\top \mathbf{h}' + b_i | \mathbf{h}']) \\
 &= \mathbb{E}_{\mathbf{h}'} [(\mathbf{h}')^\top \boldsymbol{\Sigma}_{\mathbf{W}_i} \mathbf{h}'] + \mathbb{V}_{\mathbf{h}'} (\boldsymbol{\mu}_{\mathbf{W}_i}^\top \mathbf{h}' + b_i) \\
 &= \boldsymbol{\mu}_{\mathbf{h}'}^\top \boldsymbol{\Sigma}_{\mathbf{W}_i} \boldsymbol{\mu}_{\mathbf{h}'} + \text{tr}(\boldsymbol{\Sigma}_{\mathbf{W}_i} \boldsymbol{\Sigma}_{\mathbf{h}'}) + \boldsymbol{\mu}_{\mathbf{W}_i}^\top \boldsymbol{\Sigma}_{\mathbf{h}'} \boldsymbol{\mu}_{\mathbf{W}_i} ,
 \end{aligned} \tag{8}$$

The key idea to calculate $\text{Cov}(a_1, a_2)$ is to use the law of total covariance as

$$\begin{aligned}
 \text{Cov}(a_1, a_2) &= \mathbb{E}_{\mathbf{x}} [\text{Cov}(a_1, a_2 | \mathbf{h}')] + \text{Cov}(\mathbb{E}_{\mathbf{W}} [y_1 | \mathbf{h}'], \mathbb{E}_{\mathbf{W}} [y_2 | \mathbf{h}']) \\
 &= \mathbb{E}_{\mathbf{h}'} [\text{Cov}(\mathbf{W}_1^\top \mathbf{h}' + b_1, \mathbf{W}_2^\top \mathbf{h}' + b_2 | \mathbf{h}')] + \text{Cov}(\mathbb{E}_{\mathbf{W}} [\mathbf{W}_1^\top \mathbf{h}' + b_1 | \mathbf{h}'], \mathbb{E}_{\mathbf{W}} [\mathbf{W}_2^\top \mathbf{h}' + b_2 | \mathbf{h}']) \\
 &= \mathbb{E}_{\mathbf{h}'} [\text{Cov}(\mathbf{W}_1^\top \mathbf{h}', \mathbf{W}_2^\top \mathbf{h}' | \mathbf{h}')] + \text{Cov}(\mathbb{E}_{\mathbf{W}} [\mathbf{W}_1^\top] \mathbf{h}' + b_1, \mathbb{E}_{\mathbf{W}} [\mathbf{W}_2^\top] \mathbf{h}' + b_2) \\
 &= \mathbb{E}_{\mathbf{h}'} [(\mathbf{h}')^\top \text{Cov}(\mathbf{W}_1, \mathbf{W}_2) \mathbf{h}'] + \text{Cov}(\boldsymbol{\mu}_{\mathbf{W}_1}^\top \mathbf{h}', \boldsymbol{\mu}_{\mathbf{W}_2}^\top \mathbf{h}') \\
 &= \boldsymbol{\mu}_{\mathbf{h}'}^\top \text{Cov}(\mathbf{W}_1, \mathbf{W}_2) \boldsymbol{\mu}_{\mathbf{h}'} + \text{tr}(\text{Cov}(\mathbf{W}_1, \mathbf{W}_2) \boldsymbol{\Sigma}_{\mathbf{h}'}) + \boldsymbol{\mu}_{\mathbf{W}_1}^\top \boldsymbol{\Sigma}_{\mathbf{h}'} \boldsymbol{\mu}_{\mathbf{W}_2} .
 \end{aligned} \tag{9}$$

This completes the proof. $\square$

D Full Proof of Theorem 8

We start this proof with some useful notations. We denote the Probability Density Function (PDF) of a univariate Gaussian distribution equipped with μ and σ as

$$\phi(x; \mu, \sigma^2) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(x - \mu)^2}{2\sigma^2}\right) .$$

Particularly, when $\mu = 0$ and $\sigma = 1$, this univariate Gaussian distribution is called the standard univariate Gaussian distribution, whose PDF is denoted as

$$\phi(x) = \frac{1}{\sqrt{2\pi}} \exp\left(-\frac{x^2}{2}\right) .$$

Further, we denote the Cumulative Distribution Function (CDF) of the standard univariate Gaussian distribution as

$$\Phi(x) = \int_{-\infty}^x \phi(u) du .$$

For a bivariate Gaussian distribution, we denote the PDF of a bivariate normal distribution equipped with

$$\boldsymbol{\mu} = (\mu_1, \mu_2)^\top \quad \text{and} \quad \boldsymbol{\Sigma} = \begin{pmatrix} \sigma_1^2 & \rho\sigma_1\sigma_2 \\ \rho\sigma_1\sigma_2 & \sigma_2^2 \end{pmatrix}$$

as

$$\phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) = \frac{1}{2\pi |\boldsymbol{\Sigma}|^{1/2}} \exp\left(-\frac{1}{2} (\mathbf{x} - \boldsymbol{\mu})^\top \boldsymbol{\Sigma}^{-1} (\mathbf{x} - \boldsymbol{\mu})\right) ,$$

where $\mathbf{x} = (x_1, x_2)^\top$. Especially, when

$$\boldsymbol{\mu} = (0, 0)^\top \quad \text{and} \quad \boldsymbol{\Sigma} = \begin{pmatrix} 1 & \rho \\ \rho & 1 \end{pmatrix} ,$$

this univariate Gaussian distribution is called the standard bivariate Gaussian distribution with correlation coefficient ρ , whose PDF is denoted as follows

$$\phi_2(x_1, x_2; \rho) = \frac{1}{2\pi\sqrt{1-\rho^2}} \exp\left(-\frac{1}{2(1-\rho^2)}(x_1^2 - 2\rho x_1 x_2 + x_2^2)\right).$$

Further, we denote the CDF of the standard bivariate Gaussian distribution with correlation coefficient ρ as

$$\Phi_2(x_1, x_2; \rho) = \int_{-\infty}^{x_1} \int_{-\infty}^{x_2} \phi_2(u_1, u_2; \rho) du_1 du_2.$$

With the notations above, we then introduce some useful equations as

$$\begin{aligned}\phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})|_{x_1=a} &= \phi(a; \mu_1, \sigma_1^2) \phi\left(x_2; \mu_2 - \rho \frac{\sigma_2}{\sigma_1}(\mu_1 - a), (1-\rho^2)\sigma_2^2\right), \\ \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})|_{x_2=a} &= \phi(a; \mu_2, \sigma_2^2) \phi\left(x_1; \mu_1 - \rho \frac{\sigma_1}{\sigma_2}(\mu_2 - a), (1-\rho^2)\sigma_1^2\right).\end{aligned}$$

Hence, the following holds.

$$\begin{aligned}\phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})|_{x_1=+\infty} &\doteq \lim_{a \rightarrow +\infty} \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})|_{x_1=a} = 0, \\ \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})|_{x_2=+\infty} &\doteq \lim_{a \rightarrow +\infty} \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})|_{x_2=a} = 0, \\ x_1 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})|_{x_1=0}^{x_1=+\infty} &\doteq [x_1 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})]|_{x_1=+\infty} - [x_1 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})]|_{x_1=0} = 0, \\ \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})|_{x_1=0} &= \frac{1}{\sigma_1} \phi\left(\frac{\mu_1}{\sigma_1}\right) \phi\left(x_2; \mu_2 - \rho \frac{\sigma_2}{\sigma_1} \mu_1, (1-\rho^2)\sigma_2^2\right), \\ \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})|_{x_2=0} &= \frac{1}{\sigma_2} \phi\left(\frac{\mu_2}{\sigma_2}\right) \phi\left(x_1; \mu_1 - \rho \frac{\sigma_1}{\sigma_2} \mu_2, (1-\rho^2)\sigma_1^2\right).\end{aligned}\tag{10}$$

To calculate $\mathbb{E}[h_i]$ and $\mathbb{V}(h_i)$, we follow the key idea proposed by Frey and Hinton (1999). Accordingly, it suffices to give an analytic expression on $\mathbb{E}(h_1, h_2)$, that is,

$$\begin{aligned}\mathbb{E}(h_1, h_2) &= \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \tau(x_1) \tau(x_2) \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) dx_1 dx_2 \\ &= \int_0^{+\infty} \int_0^{+\infty} x_1 x_2 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) dx_1 dx_2.\end{aligned}\tag{11}$$

The key idea is to solve the integral in Eq. (11) recursively, that is, we first build a recursive formula between the integral in Eq. (11) and $\int_0^{+\infty} \int_0^{+\infty} x_1 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) dx_1 dx_2$, and then solve the latter one. To build this recursive formula, we first note that

$$\frac{\partial \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})}{\partial \mathbf{x}} = -\boldsymbol{\Sigma}^{-1}(\mathbf{x} - \boldsymbol{\mu}) \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}).\tag{12}$$

We then multiply both sides of Eq. (12) by $\boldsymbol{\Sigma}$, explicitly write the components of the vectors $\partial \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})/\partial \mathbf{x}$ and $(\mathbf{x} - \boldsymbol{\mu})$, and obtain

$$\boldsymbol{\Sigma} \begin{pmatrix} \partial \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})/\partial x_1 \\ \partial \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})/\partial x_2 \end{pmatrix} = - \begin{pmatrix} x_1 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) - \mu_1 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) \\ x_2 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) - \mu_2 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) \end{pmatrix}.\tag{13}$$

We further unfold the left-hand side of Eq. (13) and obtain

$$\begin{pmatrix} \sigma_1^2 \partial \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})/\partial x_1 + \rho \sigma_1 \sigma_2 \partial \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})/\partial x_2 \\ \rho \sigma_1 \sigma_2 \partial \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})/\partial x_1 + \sigma_2^2 \partial \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})/\partial x_2 \end{pmatrix} = - \begin{pmatrix} x_1 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) - \mu_1 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) \\ x_2 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) - \mu_2 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) \end{pmatrix}.\tag{14}$$

Observing the first and the second components of the vectors on both sides of Eq. (14), we have

$$x_1 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) = \mu_1 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) - \sigma_1^2 \frac{\partial \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})}{\partial x_1} - \rho \sigma_1 \sigma_2 \frac{\partial \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})}{\partial x_2},\tag{15}$$

and

$$x_2 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) = \mu_2 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) - \rho \sigma_1 \sigma_2 \frac{\partial \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})}{\partial x_1} - \sigma_2^2 \frac{\partial \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})}{\partial x_2}.\tag{16}$$

Multiplying x_1 on both sides of Eq. (16), we have

$$x_1 x_2 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) = x_1 \mu_2 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) - \rho \sigma_1 \sigma_2 x_1 \frac{\partial \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})}{\partial x_1} - \sigma_2^2 x_1 \frac{\partial \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})}{\partial x_2}. \quad (17)$$

Integrating both sides of Eq. (17) over $(0, +\infty) \times (0, +\infty)$, we have the concerned recursive formula as follows.

$$\begin{aligned} \int_0^\infty \int_0^\infty x_1 x_2 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) dx_1 dx_2 &= \mu_2 \int_0^\infty \int_0^\infty x_1 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) dx_1 dx_2 - \rho \sigma_1 \sigma_2 \int_0^\infty \int_0^\infty x_1 \frac{\partial \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})}{\partial x_1} dx_1 dx_2 \\ &\quad - \sigma_2^2 \int_0^\infty \int_0^\infty x_1 \frac{\partial \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})}{\partial x_2} dx_1 dx_2 \\ &= \mu_2 \int_0^\infty \int_0^\infty x_1 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) dx_1 dx_2 - \rho \sigma_1 \sigma_2 \int_0^\infty \left(x_1 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) \Big|_{x_1=0}^{x_1=\infty} - \int_0^\infty \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) dx_1 \right) dx_2 \\ &\quad - \sigma_2^2 \int_0^\infty x_1 \left(\phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) \Big|_{x_2=0}^{x_2=\infty} \right) dx_1 \\ &= \mu_2 \int_0^\infty \int_0^\infty x_1 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) dx_1 dx_2 + \rho \sigma_1 \sigma_2 \int_0^\infty \int_0^\infty \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) dx_1 dx_2 \\ &\quad + \sigma_2 \phi\left(\frac{\mu_2}{\sigma_2}\right) \int_0^\infty x_1 \phi\left(x_1; \mu_1 - \rho \frac{\sigma_1}{\sigma_2} \mu_2, (1 - \rho^2) \sigma_1^2\right) dx_1 \\ &= \mu_2 \int_0^\infty \int_0^\infty x_1 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) dx_1 dx_2 + \rho \sigma_1 \sigma_2 \Phi_2\left(\frac{\mu_1}{\sigma_1}, \frac{\mu_2}{\sigma_2}; \rho\right) + \sigma_1 \sigma_2 \sqrt{1 - \rho^2} \phi\left(\frac{\mu_2}{\sigma_2}\right) \phi(c_1) \\ &\quad + \sigma_1 \sigma_2 \phi\left(\frac{\mu_2}{\sigma_2}\right) \left(\frac{\mu_1}{\sigma_1} - \rho \frac{\mu_2}{\sigma_2}\right) \Phi(c_1), \end{aligned} \quad (18)$$

where $c_i = (r_i - \rho r_{3-i})/\sqrt{1 - \rho^2}$ for $i \in \{1, 2\}$, the second equality uses integration by parts, and the third equality utilizes Eq. (10). Next, it suffices to solve $\int_0^{+\infty} \int_0^{+\infty} x_1 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) dx_1 dx_2$. Integrating both sides of Eq. (15) over $(0, +\infty) \times (0, +\infty)$, one has

$$\begin{aligned} \int_0^\infty \int_0^\infty x_1 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) dx_1 dx_2 &= \int_0^\infty \int_0^\infty \mu_1 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) dx_1 dx_2 - \int_0^\infty \int_0^\infty \sigma_1^2 \frac{\partial \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})}{\partial x_1} dx_1 dx_2 \\ &\quad - \int_0^\infty \int_0^\infty \rho \sigma_1 \sigma_2 \frac{\partial \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma})}{\partial x_2} dx_1 dx_2 \\ &= \mu_1 \Phi_2\left(\frac{\mu_1}{\sigma_1}, \frac{\mu_2}{\sigma_2}; \rho\right) - \sigma_1^2 \int_0^\infty \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) \Big|_{x_1=0}^{x_1=\infty} dx_2 - \rho \sigma_1 \sigma_2 \int_0^\infty \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) \Big|_{x_2=0}^{x_2=\infty} dx_1 \\ &= \mu_1 \Phi_2\left(\frac{\mu_1}{\sigma_1}, \frac{\mu_2}{\sigma_2}; \rho\right) + \sigma_1 \phi\left(\frac{\mu_1}{\sigma_1}\right) \int_0^\infty \phi\left(x_2; \mu_2 - \rho \frac{\sigma_2}{\sigma_1} \mu_1, (1 - \rho^2) \sigma_2^2\right) dx_2 \\ &\quad + \rho \sigma_1 \phi\left(\frac{\mu_2}{\sigma_2}\right) \int_0^\infty \phi\left(x_1; \mu_1 - \rho \frac{\sigma_1}{\sigma_2} \mu_2, (1 - \rho^2) \sigma_1^2\right) dx_1 \\ &= \mu_1 \Phi_2\left(\frac{\mu_1}{\sigma_1}, \frac{\mu_2}{\sigma_2}; \rho\right) + \sigma_1 \phi\left(\frac{\mu_1}{\sigma_1}\right) \Phi(c_2) + \rho \sigma_1 \phi\left(\frac{\mu_2}{\sigma_2}\right) \Phi(c_1), \end{aligned} \quad (19)$$

where $r_i = \mu_i/\sigma_i$ for $i \in \{1, 2\}$, the second equality utilizes integration by parts, and the third equality uses Eq. (10). Substituting Eq. (19) into Eq. (18), one has

$$\begin{aligned} \int_0^\infty \int_0^\infty x_1 x_2 \phi_2(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) dx_1 dx_2 &= \\ &\quad \sigma_1 \sigma_2 \left[r_2 \phi(r_1) \Phi(c_2) + r_1 \phi(r_2) \Phi(c_1) + (r_1 r_2 + \rho) \Phi_2(r_1, r_2; \rho) + \sqrt{1 - \rho^2} \phi(r_2) \phi(c_1) \right]. \end{aligned}$$

This completes the proof. $\square$

E Full Proof of Theorem 9

According to Corollary 5, if

$$\begin{cases} q^{(l)} - p^{(l)} = 0, \\ p^{(l)} + q^{(l)} - \left(q^{(l)} - p^{(l)}\right)^2 = 0, \end{cases}$$

for $l \in [L]$, then one has

$$\mathbb{E}[\mathbf{a}^{(l)}] = \mathbf{0}, \quad \mathbb{V}[\mathbf{a}^{(l)}] = \mathbf{0}, \quad \text{Cov}(\mathbf{a}^{(l)}, \mathbf{a}^{(l)}) = \mathbf{0},$$

where

$$\mathbf{a}^{(l)} = \mathbf{W}^{(l)} \mathbf{h}^{(l-1)} + \mathbf{b}^{(l)}.$$

It is obvious that

$$\Sigma_{\hat{\mathbf{y}}} = \text{Cov}(\mathbf{a}^{(L)}, \mathbf{a}^{(L)}) = \mathbf{0},$$

which completes the proof. $\square$

This theorem admits an alternative proof. It suffices to note that

$$\begin{cases} q^{(l)} - p^{(l)} = 0, \\ p^{(l)} + q^{(l)} - (q^{(l)} - p^{(l)})^2 = 0, \end{cases}$$

are equivalent to

$$p^{(l)} = q^{(l)} = 0 \quad \text{for } l \in [L].$$

Thus, distribution $\mathcal{T}(p^{(l)}, 1 - p^{(l)} - q^{(l)}, q^{(l)})$ is considered deterministic, leading to $\Sigma_{\hat{\mathbf{y}}} = \mathbf{0}$ and thus completing the proof. $\square$

F Experiments Details of Subsection 5.1

This appendix provides the experimental details for Subsection 5.1.

Table 3, Table 4, and Table 5 show the uncertainty estimation performance of EMP and its contenders, including PL-DNN (Bibi, Alfadly, and Ghanem 2018) and MP (Wu et al. 2019) across datasets for MLP-2, LeNet-5, and VGG-13, respectively. We plot the results of only the first 10 components of CIFAR-100 due to space constraints. It is observed that the EMP method achieves near-ideal performance and significantly outperforms its contenders. This observation supports the effectiveness of the proposed EMP estimator.

Datasets	o	FP-MLP-2			T-MLP-2		
		EMP	MP	PL-DNN	EMP	MP	PL-DNN
MNIST	0	1.0001 _{0.0002} $\pm$ 0.0101 _{0.0008}	1.0818 _{0.0226} $\pm$ 0.2029 _{0.0401}	0.9083 _{0.0042} $\pm$ 0.1957 _{0.0250}	0.9954 _{0.0004} $\pm$ 0.0103 _{0.0010}	—	—
	1	1.0000 _{0.0001} $\pm$ 0.0100 _{0.0008}	1.0818 _{0.0226} $\pm$ 0.2032 _{0.0400}	0.9083 _{0.0043} $\pm$ 0.1957 _{0.0250}	0.9953 _{0.0003} $\pm$ 0.0104 _{0.0010}	—	—
	2	0.9999 _{0.0001} $\pm$ 0.0100 _{0.0008}	1.0818 _{0.0226} $\pm$ 0.2031 _{0.0401}	0.9086 _{0.0044} $\pm$ 0.1958 _{0.0250}	0.9953 _{0.0003} $\pm$ 0.0104 _{0.0009}	—	—
	3	0.9998 _{0.0002} $\pm$ 0.0100 _{0.0007}	1.0818 _{0.0226} $\pm$ 0.2029 _{0.0401}	0.9085 _{0.0043} $\pm$ 0.1958 _{0.0250}	0.9955 _{0.0002} $\pm$ 0.0105 _{0.0009}	—	—
	4	1.0000 _{0.0002} $\pm$ 0.0099 _{0.0008}	1.0817 _{0.0225} $\pm$ 0.2030 _{0.0401}	0.9084 _{0.0044} $\pm$ 0.1957 _{0.0250}	0.9956 _{0.0004} $\pm$ 0.0103 _{0.0010}	—	—
	5	1.0000 _{0.0001} $\pm$ 0.0100 _{0.0008}	1.0818 _{0.0225} $\pm$ 0.2031 _{0.0401}	0.9084 _{0.0044} $\pm$ 0.1957 _{0.0250}	0.9954 _{0.0004} $\pm$ 0.0104 _{0.0010}	—	—
	6	1.0000 _{0.0000} $\pm$ 0.0099 _{0.0007}	1.0817 _{0.0226} $\pm$ 0.2030 _{0.0399}	0.9084 _{0.0043} $\pm$ 0.1959 _{0.0250}	0.9952 _{0.0005} $\pm$ 0.0104 _{0.0009}	—	—
	7	1.0001 _{0.0001} $\pm$ 0.0100 _{0.0007}	1.0818 _{0.0226} $\pm$ 0.2032 _{0.0400}	0.9084 _{0.0043} $\pm$ 0.1957 _{0.0250}	0.9952 _{0.0005} $\pm$ 0.0104 _{0.0009}	—	—
	8	0.9999 _{0.0002} $\pm$ 0.0101 _{0.0007}	1.0819 _{0.0224} $\pm$ 0.2030 _{0.0402}	0.9083 _{0.0044} $\pm$ 0.1956 _{0.0250}	0.9951 _{0.0004} $\pm$ 0.0106 _{0.0009}	—	—
	9	1.0000 _{0.0001} $\pm$ 0.0100 _{0.0007}	1.0819 _{0.0226} $\pm$ 0.2031 _{0.0400}	0.9084 _{0.0043} $\pm$ 0.1959 _{0.0249}	0.9954 _{0.0004} $\pm$ 0.0104 _{0.0010}	—	—
CIFAR-10	0	1.0000 _{0.0003} $\pm$ 0.0077 _{0.0009}	1.1916 _{0.0281} $\pm$ 0.2538 _{0.0167}	0.7905 _{0.0258} $\pm$ 0.3391 _{0.0212}	0.9947 _{0.0007} $\pm$ 0.0108 _{0.0011}	—	—
	1	0.9999 _{0.0003} $\pm$ 0.0078 _{0.0011}	1.1912 _{0.0275} $\pm$ 0.2535 _{0.0165}	0.7903 _{0.0263} $\pm$ 0.3389 _{0.0212}	0.9947 _{0.0006} $\pm$ 0.0106 _{0.0013}	—	—
	2	1.0001 _{0.0003} $\pm$ 0.0077 _{0.0011}	1.1912 _{0.0276} $\pm$ 0.2537 _{0.0165}	0.7905 _{0.0264} $\pm$ 0.3390 _{0.0213}	0.9948 _{0.0006} $\pm$ 0.0109 _{0.0012}	—	—
	3	0.9999 _{0.0002} $\pm$ 0.0077 _{0.0011}	1.1916 _{0.0277} $\pm$ 0.2537 _{0.0165}	0.7906 _{0.0261} $\pm$ 0.3391 _{0.0212}	0.9945 _{0.0005} $\pm$ 0.0108 _{0.0013}	—	—
	4	0.9998 _{0.0002} $\pm$ 0.0078 _{0.0012}	1.1914 _{0.0279} $\pm$ 0.2538 _{0.0164}	0.7905 _{0.0263} $\pm$ 0.3392 _{0.0211}	0.9947 _{0.0007} $\pm$ 0.0106 _{0.0012}	—	—
	5	1.0001 _{0.0003} $\pm$ 0.0076 _{0.0010}	1.1913 _{0.0276} $\pm$ 0.2539 _{0.0169}	0.7905 _{0.0262} $\pm$ 0.3390 _{0.0213}	0.9948 _{0.0007} $\pm$ 0.0107 _{0.0011}	—	—
	6	1.0000 _{0.0003} $\pm$ 0.0078 _{0.0008}	1.1915 _{0.0279} $\pm$ 0.2538 _{0.0166}	0.7905 _{0.0263} $\pm$ 0.3392 _{0.0210}	0.9946 _{0.0008} $\pm$ 0.0107 _{0.0015}	—	—
	7	1.0003 _{0.0002} $\pm$ 0.0077 _{0.0011}	1.1916 _{0.0275} $\pm$ 0.2539 _{0.0165}	0.7907 _{0.0260} $\pm$ 0.3392 _{0.0211}	0.9946 _{0.0008} $\pm$ 0.0108 _{0.0015}	—	—
	8	0.9999 _{0.0003} $\pm$ 0.0079 _{0.0012}	1.1913 _{0.0277} $\pm$ 0.2536 _{0.0165}	0.7906 _{0.0263} $\pm$ 0.3392 _{0.0212}	0.9946 _{0.0007} $\pm$ 0.0107 _{0.0011}	—	—
	9	0.9999 _{0.0001} $\pm$ 0.0078 _{0.0011}	1.1914 _{0.0275} $\pm$ 0.2537 _{0.0166}	0.7904 _{0.0263} $\pm$ 0.3391 _{0.0211}	0.9948 _{0.0007} $\pm$ 0.0105 _{0.0012}	—	—
CIFAR-100	0	1.0003 _{0.0008} $\pm$ 0.0074 _{0.0008}	1.2432 _{0.0339} $\pm$ 0.2675 _{0.0154}	0.6636 _{0.0217} $\pm$ 0.4170 _{0.0146}	0.9942 _{0.0012} $\pm$ 0.0104 _{0.0014}	—	—
	1	1.0004 _{0.0007} $\pm$ 0.0070 _{0.0006}	1.2430 _{0.0341} $\pm$ 0.2673 _{0.0153}	0.6630 _{0.0214} $\pm$ 0.4170 _{0.0146}	0.9946 _{0.0016} $\pm$ 0.0105 _{0.0012}	—	—
	2	1.0001 _{0.0006} $\pm$ 0.0072 _{0.0008}	1.2424 _{0.0334} $\pm$ 0.2661 _{0.0145}	0.6634 _{0.0215} $\pm$ 0.4170 _{0.0145}	0.9946 _{0.0015} $\pm$ 0.0105 _{0.0011}	—	—
	3	0.9997 _{0.0007} $\pm$ 0.0066 _{0.0011}	1.2428 _{0.0339} $\pm$ 0.2673 _{0.0150}	0.6631 _{0.0210} $\pm$ 0.4174 _{0.0146}	0.9948 _{0.0015} $\pm$ 0.0108 _{0.0011}	—	—
	4	1.0003 _{0.0010} $\pm$ 0.0070 _{0.0012}	1.2436 _{0.0337} $\pm$ 0.2674 _{0.0154}	0.6635 _{0.0212} $\pm$ 0.4170 _{0.0150}	0.9946 _{0.0015} $\pm$ 0.0108 _{0.0015}	—	—
	5	1.0001 _{0.0004} $\pm$ 0.0069 _{0.0008}	1.2433 _{0.0343} $\pm$ 0.2673 _{0.0150}	0.6632 _{0.0207} $\pm$ 0.4169 _{0.0144}	0.9942 _{0.0011} $\pm$ 0.0107 _{0.0012}	—	—
	6	1.0000 _{0.0003} $\pm$ 0.0069 _{0.0011}	1.2436 _{0.0338} $\pm$ 0.2676 _{0.0149}	0.6630 _{0.0211} $\pm$ 0.4168 _{0.0149}	0.9943 _{0.0016} $\pm$ 0.0105 _{0.0012}	—	—
	7	1.0002 _{0.0006} $\pm$ 0.0073 _{0.0008}	1.2429 _{0.0337} $\pm$ 0.2668 _{0.0150}	0.6633 _{0.0217} $\pm$ 0.4171 _{0.0148}	0.9943 _{0.0013} $\pm$ 0.0107 _{0.0014}	—	—
	8	1.0000 _{0.0006} $\pm$ 0.0066 _{0.0006}	1.2431 _{0.0343} $\pm$ 0.2675 _{0.0154}	0.6632 _{0.0214} $\pm$ 0.4169 _{0.0148}	0.9943 _{0.0017} $\pm$ 0.0107 _{0.0012}	—	—
	9	1.0000 _{0.0006} $\pm$ 0.0071 _{0.0012}	1.2431 _{0.0339} $\pm$ 0.2675 _{0.0158}	0.6632 _{0.0212} $\pm$ 0.4168 _{0.0146}	0.9948 _{0.0013} $\pm$ 0.0106 _{0.0015}	—	—

Table 3: Performance of the EMP and contenders for MLP-2 across datasets.

Datasets	σ	FP-MLP-2			T-MLP-2		
		EMP	MP	PL-DNN	EMP	MP	PL-DNN
MNIST	0	1.0007 _{0.0006} $\pm$ 0.0101 _{0.0002}	1.1809 _{0.0301} $\pm$ 0.0217 _{0.0069}	1.2772 _{0.0497} $\pm$ 0.0153 _{0.0011}	1.0023 _{0.0112} $\pm$ 0.0109 _{0.0017}	—	—
	1	0.9999 _{0.0002} $\pm$ 0.0103 _{0.0000}	1.1793 _{0.0304} $\pm$ 0.0238 _{0.0058}	1.2759 _{0.0499} $\pm$ 0.0139 _{0.0008}	1.0023 _{0.0111} $\pm$ 0.0121 _{0.0012}	—	—
	2	0.9990 _{0.0005} $\pm$ 0.0118 _{0.0007}	1.1794 _{0.0297} $\pm$ 0.0239 _{0.0060}	1.2748 _{0.0501} $\pm$ 0.0170 _{0.0015}	1.0021 _{0.0110} $\pm$ 0.0117 _{0.0011}	—	—
	3	1.0008 _{0.0004} $\pm$ 0.0104 _{0.0001}	1.1808 _{0.0301} $\pm$ 0.0215 _{0.0065}	1.2770 _{0.0498} $\pm$ 0.0153 _{0.0011}	1.0022 _{0.0110} $\pm$ 0.0108 _{0.0016}	—	—
	4	1.0011 _{0.0006} $\pm$ 0.0096 _{0.0004}	1.1807 _{0.0304} $\pm$ 0.0239 _{0.0067}	1.2775 _{0.0497} $\pm$ 0.0136 _{0.0011}	1.0023 _{0.0109} $\pm$ 0.0107 _{0.0016}	—	—
	5	1.0010 _{0.0006} $\pm$ 0.0106 _{0.0001}	1.1815 _{0.0300} $\pm$ 0.0237 _{0.0059}	1.2773 _{0.0498} $\pm$ 0.0149 _{0.0010}	1.0022 _{0.0109} $\pm$ 0.0109 _{0.0013}	—	—
	6	1.0011 _{0.0007} $\pm$ 0.0096 _{0.0004}	1.1818 _{0.0299} $\pm$ 0.0246 _{0.0062}	1.2774 _{0.0496} $\pm$ 0.0130 _{0.0011}	1.0023 _{0.0112} $\pm$ 0.0105 _{0.0013}	—	—
	7	0.9996 _{0.0002} $\pm$ 0.0099 _{0.0002}	1.1803 _{0.0296} $\pm$ 0.0243 _{0.0067}	1.2755 _{0.0500} $\pm$ 0.0143 _{0.0009}	1.0021 _{0.0109} $\pm$ 0.0101 _{0.0011}	—	—
	8	1.0014 _{0.0007} $\pm$ 0.0098 _{0.0003}	1.1814 _{0.0301} $\pm$ 0.0218 _{0.0066}	1.2422 _{0.0714} $\pm$ 0.0140 _{0.0010}	1.0021 _{0.0111} $\pm$ 0.0107 _{0.0012}	—	—
	9	0.9999 _{0.0004} $\pm$ 0.0100 _{0.0001}	1.1801 _{0.0298} $\pm$ 0.0225 _{0.0063}	1.2759 _{0.0497} $\pm$ 0.0142 _{0.0009}	1.0019 _{0.0111} $\pm$ 0.0106 _{0.0013}	—	—
CIFAR-10	0	1.0098 _{0.0017} $\pm$ 0.0099 _{0.0001}	1.5029 _{0.0445} $\pm$ 0.0177 _{0.0010}	1.6167 _{0.0559} $\pm$ 0.0334 _{0.0022}	1.0141 _{0.0133} $\pm$ 0.0113 _{0.0017}	—	—
	1	1.0092 _{0.0017} $\pm$ 0.0105 _{0.0001}	1.5020 _{0.0446} $\pm$ 0.0207 _{0.0011}	1.6156 _{0.0559} $\pm$ 0.0299 _{0.0022}	1.0136 _{0.0133} $\pm$ 0.0115 _{0.0018}	—	—
	2	1.0090 _{0.0018} $\pm$ 0.0127 _{0.0002}	1.5017 _{0.0447} $\pm$ 0.0218 _{0.0009}	1.6154 _{0.0560} $\pm$ 0.0358 _{0.0022}	1.0141 _{0.0135} $\pm$ 0.0110 _{0.0017}	—	—
	3	1.0100 _{0.0017} $\pm$ 0.0110 _{0.0002}	1.5031 _{0.0446} $\pm$ 0.0189 _{0.0010}	1.6169 _{0.0558} $\pm$ 0.0338 _{0.0023}	1.0141 _{0.0133} $\pm$ 0.0115 _{0.0021}	—	—
	4	1.0106 _{0.0017} $\pm$ 0.0098 _{0.0004}	1.5041 _{0.0446} $\pm$ 0.0179 _{0.0011}	1.6179 _{0.0557} $\pm$ 0.0307 _{0.0023}	1.0137 _{0.0136} $\pm$ 0.0116 _{0.0015}	—	—
	5	1.0099 _{0.0017} $\pm$ 0.0109 _{0.0001}	1.5030 _{0.0445} $\pm$ 0.0200 _{0.0011}	1.6167 _{0.0556} $\pm$ 0.0311 _{0.0022}	1.0139 _{0.0133} $\pm$ 0.0118 _{0.0022}	—	—
	6	1.0103 _{0.0017} $\pm$ 0.0094 _{0.0002}	1.5036 _{0.0445} $\pm$ 0.0185 _{0.0011}	1.6174 _{0.0556} $\pm$ 0.0282 _{0.0022}	1.0139 _{0.0134} $\pm$ 0.0109 _{0.0018}	—	—
	7	1.0090 _{0.0017} $\pm$ 0.0099 _{0.0003}	1.5017 _{0.0446} $\pm$ 0.0184 _{0.0010}	1.6153 _{0.0559} $\pm$ 0.0307 _{0.0022}	1.0139 _{0.0134} $\pm$ 0.0120 _{0.0014}	—	—
	8	1.0103 _{0.0017} $\pm$ 0.0093 _{0.0001}	1.5036 _{0.0445} $\pm$ 0.0179 _{0.0010}	1.6173 _{0.0558} $\pm$ 0.0312 _{0.0022}	1.0140 _{0.0133} $\pm$ 0.0118 _{0.0015}	—	—
	9	1.0095 _{0.0017} $\pm$ 0.0100 _{0.0001}	1.5025 _{0.0446} $\pm$ 0.0179 _{0.0010}	1.6162 _{0.0559} $\pm$ 0.0314 _{0.0023}	1.0141 _{0.0135} $\pm$ 0.0109 _{0.0020}	—	—
CIFAR-100	0	1.0133 _{0.0006} $\pm$ 0.0128 _{0.0002}	0.6602 _{0.0171} $\pm$ 0.0959 _{0.0034}	0.5738 _{0.0204} $\pm$ 0.1323 _{0.0040}	1.0135 _{0.0131} $\pm$ 0.0112 _{0.0016}	—	—
	1	1.0134 _{0.0006} $\pm$ 0.0134 _{0.0001}	0.6603 _{0.0171} $\pm$ 0.0961 _{0.0033}	0.5740 _{0.0205} $\pm$ 0.1324 _{0.0039}	1.0132 _{0.0128} $\pm$ 0.0117 _{0.0020}	—	—
	2	1.0135 _{0.0006} $\pm$ 0.0131 _{0.0002}	0.6602 _{0.0171} $\pm$ 0.0949 _{0.0034}	0.5738 _{0.0205} $\pm$ 0.1314 _{0.0040}	1.0133 _{0.0128} $\pm$ 0.0114 _{0.0018}	—	—
	3	1.0122 _{0.0006} $\pm$ 0.0106 _{0.0001}	0.6598 _{0.0170} $\pm$ 0.0974 _{0.0033}	0.5736 _{0.0204} $\pm$ 0.1336 _{0.0039}	1.0136 _{0.0129} $\pm$ 0.0114 _{0.0017}	—	—
	4	1.0134 _{0.0007} $\pm$ 0.0121 _{0.0002}	0.6603 _{0.0171} $\pm$ 0.0960 _{0.0034}	0.5740 _{0.0205} $\pm$ 0.1324 _{0.0040}	1.0132 _{0.0129} $\pm$ 0.0112 _{0.0012}	—	—
	5	1.0126 _{0.0006} $\pm$ 0.0120 _{0.0002}	0.6598 _{0.0171} $\pm$ 0.0959 _{0.0033}	0.5736 _{0.0204} $\pm$ 0.1323 _{0.0040}	1.0132 _{0.0128} $\pm$ 0.0118 _{0.0015}	—	—
	6	1.0125 _{0.0007} $\pm$ 0.0135 _{0.0002}	0.6596 _{0.0172} $\pm$ 0.0951 _{0.0034}	0.5732 _{0.0205} $\pm$ 0.1316 _{0.0040}	1.0134 _{0.0130} $\pm$ 0.0117 _{0.0019}	—	—
	7	1.0126 _{0.0006} $\pm$ 0.0121 _{0.0002}	0.6600 _{0.0171} $\pm$ 0.0970 _{0.0034}	0.5737 _{0.0204} $\pm$ 0.1333 _{0.0040}	1.0134 _{0.0128} $\pm$ 0.0118 _{0.0019}	—	—
	8	1.0133 _{0.0006} $\pm$ 0.0116 _{0.0001}	0.6603 _{0.0171} $\pm$ 0.0964 _{0.0033}	0.5739 _{0.0205} $\pm$ 0.1328 _{0.0039}	1.0138 _{0.0128} $\pm$ 0.0110 _{0.0019}	—	—
	9	1.0131 _{0.0006} $\pm$ 0.0118 _{0.0001}	0.6602 _{0.0171} $\pm$ 0.0963 _{0.0033}	0.5739 _{0.0205} $\pm$ 0.1327 _{0.0040}	1.0133 _{0.0127} $\pm$ 0.0112 _{0.0017}	—	—

Table 4: Performance of the EMP and contenders for LeNet-5 across datasets.

Datasets	σ	FP-MLP-2			T-MLP-2		
		EMP	MP	PL-DNN	EMP	MP	PL-DNN
MNIST	0	0.8923 _{0.0521} ± 0.0089 _{0.0017}	0.5461 _{0.1998} ± 0.0736 _{0.0052}	0.5420 _{0.2202} ± 0.0741 _{0.0054}	0.8875 _{0.0633} ± 0.0174 _{0.0013}	—	—
	1	0.8934 _{0.0520} ± 0.0089 _{0.0016}	0.5476 _{0.1993} ± 0.0688 _{0.0048}	0.5435 _{0.2297} ± 0.0694 _{0.0049}	0.8873 _{0.0635} ± 0.0177 _{0.0015}	—	—
	2	0.8937 _{0.0523} ± 0.0091 _{0.0018}	0.5483 _{0.1995} ± 0.0712 _{0.0050}	0.5442 _{0.2299} ± 0.0718 _{0.0051}	0.8869 _{0.0632} ± 0.0169 _{0.0014}	—	—
	3	0.8935 _{0.0522} ± 0.0093 _{0.0017}	0.5489 _{0.1996} ± 0.0722 _{0.0053}	0.5448 _{0.2201} ± 0.0727 _{0.0054}	0.8873 _{0.0634} ± 0.0174 _{0.0012}	—	—
	4	0.8935 _{0.0521} ± 0.0089 _{0.0016}	0.5490 _{0.1997} ± 0.0727 _{0.0055}	0.5449 _{0.2200} ± 0.0732 _{0.0056}	0.8875 _{0.0631} ± 0.0172 _{0.0014}	—	—
	5	0.8935 _{0.0520} ± 0.0089 _{0.0017}	0.5492 _{0.1994} ± 0.0704 _{0.0051}	0.5451 _{0.2298} ± 0.0710 _{0.0052}	0.8868 _{0.0635} ± 0.0173 _{0.0013}	—	—
	6	0.8940 _{0.0524} ± 0.0091 _{0.0018}	0.5495 _{0.1992} ± 0.0702 _{0.0049}	0.5454 _{0.2296} ± 0.0707 _{0.0050}	0.8874 _{0.0633} ± 0.0176 _{0.0015}	—	—
	7	0.8934 _{0.0522} ± 0.0090 _{0.0017}	0.5483 _{0.1995} ± 0.0713 _{0.0054}	0.5442 _{0.2299} ± 0.0718 _{0.0055}	0.8867 _{0.0634} ± 0.0177 _{0.0013}	—	—
	8	0.8934 _{0.0519} ± 0.0087 _{0.0016}	0.5504 _{0.1999} ± 0.0734 _{0.0057}	0.5462 _{0.2203} ± 0.0739 _{0.0058}	0.8874 _{0.0632} ± 0.0173 _{0.0014}	—	—
	9	0.8937 _{0.0521} ± 0.0089 _{0.0017}	0.5505 _{0.1993} ± 0.0706 _{0.0052}	0.5463 _{0.2297} ± 0.0712 _{0.0053}	0.8876 _{0.0631} ± 0.0172 _{0.0012}	—	—
CIFAR-10	0	0.8363 _{0.0625} ± 0.0153 _{0.0021}	0.3579 _{0.1534} ± 0.1387 _{0.0075}	0.3567 _{0.1535} ± 0.1389 _{0.0076}	0.8441 _{0.0757} ± 0.0382 _{0.0028}	—	—
	1	0.8358 _{0.0624} ± 0.0152 _{0.0020}	0.3577 _{0.1532} ± 0.1380 _{0.0073}	0.3565 _{0.1533} ± 0.1381 _{0.0074}	0.8434 _{0.0755} ± 0.0376 _{0.0026}	—	—
	2	0.8361 _{0.0623} ± 0.0146 _{0.0019}	0.3596 _{0.1539} ± 0.1446 _{0.0078}	0.3584 _{0.1540} ± 0.1447 _{0.0079}	0.8442 _{0.0758} ± 0.0385 _{0.0029}	—	—
	3	0.8361 _{0.0626} ± 0.0158 _{0.0022}	0.3569 _{0.1535} ± 0.1392 _{0.0076}	0.3557 _{0.1536} ± 0.1393 _{0.0077}	0.8442 _{0.0757} ± 0.0385 _{0.0028}	—	—
	4	0.8363 _{0.0626} ± 0.0158 _{0.0021}	0.3562 _{0.1534} ± 0.1389 _{0.0075}	0.3550 _{0.1535} ± 0.1390 _{0.0076}	0.8441 _{0.0756} ± 0.0383 _{0.0027}	—	—
	5	0.8359 _{0.0624} ± 0.0148 _{0.0020}	0.3560 _{0.1531} ± 0.1377 _{0.0072}	0.3548 _{0.1532} ± 0.1378 _{0.0073}	0.8441 _{0.0756} ± 0.0381 _{0.0026}	—	—
	6	0.8360 _{0.0624} ± 0.0149 _{0.0019}	0.3571 _{0.1538} ± 0.1418 _{0.0077}	0.3559 _{0.1539} ± 0.1419 _{0.0078}	0.8444 _{0.0754} ± 0.0375 _{0.0025}	—	—
	7	0.8348 _{0.0625} ± 0.0152 _{0.0021}	0.3575 _{0.1534} ± 0.1390 _{0.0076}	0.3563 _{0.1535} ± 0.1391 _{0.0077}	0.8435 _{0.0757} ± 0.0382 _{0.0028}	—	—
	8	0.8356 _{0.0623} ± 0.0147 _{0.0019}	0.3578 _{0.1536} ± 0.1401 _{0.0078}	0.3566 _{0.1537} ± 0.1403 _{0.0079}	0.8433 _{0.0757} ± 0.0383 _{0.0027}	—	—
	9	0.8362 _{0.0625} ± 0.0151 _{0.0020}	0.3565 _{0.1530} ± 0.1366 _{0.0071}	0.3553 _{0.1531} ± 0.1367 _{0.0072}	0.8443 _{0.0757} ± 0.0384 _{0.0029}	—	—
CIFAR-100	0	0.8258 _{0.0636} ± 0.0246 _{0.0028}	0.3926 _{0.1439} ± 0.1445 _{0.0082}	0.3914 _{0.1580} ± 0.1447 _{0.0083}	0.8553 _{0.0748} ± 0.0327 _{0.0023}	—	—
	1	0.8256 _{0.0636} ± 0.0246 _{0.0027}	0.3922 _{0.1438} ± 0.1439 _{0.0081}	0.3909 _{0.1589} ± 0.1440 _{0.0082}	0.8544 _{0.0746} ± 0.0321 _{0.0022}	—	—
	2	0.8256 _{0.0638} ± 0.0253 _{0.0029}	0.3943 _{0.1444} ± 0.1508 _{0.0085}	0.3931 _{0.1585} ± 0.1510 _{0.0086}	0.8554 _{0.0748} ± 0.0328 _{0.0024}	—	—
	3	0.8251 _{0.0637} ± 0.0247 _{0.0028}	0.3914 _{0.1440} ± 0.1451 _{0.0083}	0.3902 _{0.1581} ± 0.1453 _{0.0084}	0.8552 _{0.0746} ± 0.0321 _{0.0021}	—	—
	4	0.8251 _{0.0637} ± 0.0247 _{0.0028}	0.3906 _{0.1439} ± 0.1448 _{0.0082}	0.3894 _{0.1580} ± 0.1450 _{0.0083}	0.8556 _{0.0749} ± 0.0328 _{0.0024}	—	—
	5	0.8249 _{0.0638} ± 0.0252 _{0.0029}	0.3905 _{0.1438} ± 0.1434 _{0.0080}	0.3892 _{0.1589} ± 0.1436 _{0.0081}	0.8548 _{0.0747} ± 0.0326 _{0.0023}	—	—
	6	0.8258 _{0.0638} ± 0.0251 _{0.0029}	0.3915 _{0.1442} ± 0.1478 _{0.0084}	0.3902 _{0.1583} ± 0.1479 _{0.0085}	0.8546 _{0.0746} ± 0.0322 _{0.0022}	—	—
	7	0.8250 _{0.0637} ± 0.0248 _{0.0028}	0.3920 _{0.1439} ± 0.1449 _{0.0082}	0.3908 _{0.1580} ± 0.1451 _{0.0083}	0.8552 _{0.0748} ± 0.0326 _{0.0023}	—	—
	8	0.8251 _{0.0637} ± 0.0247 _{0.0028}	0.3924 _{0.1441} ± 0.1460 _{0.0083}	0.3912 _{0.1582} ± 0.1462 _{0.0084}	0.8553 _{0.0748} ± 0.0327 _{0.0023}	—	—
	9	0.8253 _{0.0638} ± 0.0253 _{0.0029}	0.3910 _{0.1437} ± 0.1424 _{0.0079}	0.3898 _{0.1588} ± 0.1426 _{0.0080}	0.8554 _{0.0749} ± 0.0328 _{0.0024}	—	—

G Experiments Details of Subsection 5.2

This appendix provides the experimental details for Subsection 5.2.

The hyperparameters of Algorithm 2 are chosen as $c = 0.3$ and $\lambda = 10$ for the MLP-2 architecture, and set as $c = 0.9$ and $\lambda = 1$ for other architectures.

Table 6 shows the uncertainty reduction performance of MOMA by comparing the uncertainty of the T-DNNs and T-DNNs-MOMA. It is observed that the MOMA greatly reduces the uncertainty of T-DNNs across various datasets and DNN architectures. This observation supports the effectiveness of the proposed MOMA method.

Datasets	Models	Accuracy	Uncertainty	Reduction
MNIST	FP-MLP-2	0.9761 _{0.0002}	$1.4697_{0.0082} \times 10^3$	68.08 _{0.1389} %
	T-MLP-2	0.9752 _{0.0005}	$1.2572_{0.0236} \times 10^5$	
	T-MLP-2-MOMA	0.9751 _{0.0006}	$4.0129_{0.0039} \times 10^4$	
	FP-LeNet-5	0.9911 _{0.0001}	$1.7728_{0.0232} \times 10^3$	98.55 _{0.0025} %
	T-LeNet-5	0.9903 _{0.0003}	$2.8113_{0.0490} \times 10^5$	
	T-LeNet-5-MOMA	0.9885 _{0.0006}	$4.0906_{0.1433} \times 10^3$	
	FP-VGG-13	0.9985 _{0.0001}	$4.6029_{0.2302} \times 10^{-1}$	20.86 _{2.5366} %
	T-VGG-13	0.9954 _{0.0005}	$1.4500_{0.0210} \times 10^1$	
	T-VGG-13-MOMA	0.9957 _{0.0005}	$1.1463_{0.0198} \times 10^1$	
CIFAR-10	FP-ResNet-18	0.9963 _{0.0001}	$6.9757_{0.3743} \times 10^{-3}$	44.27 _{4.4212} %
	T-ResNet-18	0.9938 _{0.0002}	$4.2808_{0.1772} \times 10^{-1}$	
	T-ResNet-18-MOMA	0.9951 _{0.0002}	$2.4067_{0.0501} \times 10^{-1}$	
	FP-MLP-2	0.4806 _{0.0010}	$5.8953_{0.1782} \times 10^3$	86.20 _{1.0919} %
	T-MLP-2	0.1066 _{0.0058}	$5.6167_{0.3858} \times 10^6$	
	T-MLP-2-MOMA	0.1055 _{0.0059}	$7.7525_{0.3044} \times 10^5$	
	FP-LeNet-5	0.6651 _{0.0003}	$2.0605_{0.2572} \times 10^{-1}$	98.53 _{1.5265} %
	T-LeNet-5	0.6140 _{0.0062}	$4.3177_{2.0793} \times 10^2$	
	T-LeNet-5-MOMA	0.6138 _{0.0066}	$6.3389_{5.8414} \times 10^0$	
CIFAR-100	FP-VGG-13	0.9263 _{0.0002}	$3.1242_{0.3023} \times 10^{-1}$	31.67 _{4.0737} %
	T-VGG-13	0.9212 _{0.0082}	$9.9900_{0.4232} \times 10^1$	
	T-VGG-13-MOMA	0.9096 _{0.0083}	$6.7903_{0.3624} \times 10^1$	
	FP-ResNet-18	0.9395 _{0.0002}	$4.4530_{0.3822} \times 10^{-2}$	18.27 _{0.9785} %
	T-ResNet-18	0.9174 _{0.0076}	$2.9628_{0.0252} \times 10^{-1}$	
	T-ResNet-18-MOMA	0.9233 _{0.0075}	$2.4215_{0.0249} \times 10^{-1}$	
	FP-MLP-2	0.0794 _{0.0015}	$1.0644_{0.0472} \times 10^5$	69.90 _{0.7244} %
	T-MLP-2	0.0624 _{0.0066}	$9.9100_{0.1074} \times 10^6$	
	T-MLP-2-MOMA	0.0725 _{0.0022}	$2.9831_{0.0641} \times 10^6$	
CIFAR-100	FP-LeNet-5	0.2950 _{0.0024}	$3.4251_{0.0548} \times 10^3$	30.93 _{0.5204} %
	T-LeNet-5	0.0102 _{0.0001}	$1.0975_{1.7877} \times 10^5$	
	T-LeNet-5-MOMA	0.0101 _{0.0001}	$7.5830_{0.4088} \times 10^4$	
	FP-VGG-13	0.7026 _{0.0009}	$3.6694_{0.0772} \times 10^{-1}$	12.51 _{5.4372} %
	T-VGG-13	0.6806 _{0.0045}	$2.1427_{0.1035} \times 10^1$	
	T-VGG-13-MOMA	0.6705 _{0.0042}	$1.8700_{0.0981} \times 10^1$	
	FP-ResNet-18	0.7354 _{0.0007}	$1.6989_{0.2109} \times 10^{-2}$	16.14 _{5.6127} %
	T-ResNet-18	0.6672 _{0.0038}	$1.5869_{0.0811} \times 10^{-1}$	
	T-ResNet-18-MOMA	0.6880 _{0.0040}	$1.3263_{0.0762} \times 10^{-1}$	

Table 6: Accuracy and uncertainty of concerned models on MNIST, CIFAR-10, and CIFAR-100.

Figure 5, Figure 6, Figure 7, and Figure 8 are results for the MLP-2 on MNIST across different ratio ranges, and they serve two purposes. Firstly, they show the uncertainty reduction performance of MOMA by visualizing the uncertainty ratio of samples across a dataset. It is observed that the red bars are taller at lower x-values and shorter at higher x-values. This observation supports the effectiveness of the proposed MOMA method. Secondly, they show that weight ternarization enlarges DNN uncertainty, regardless of whether the MOMA is used during WQ.

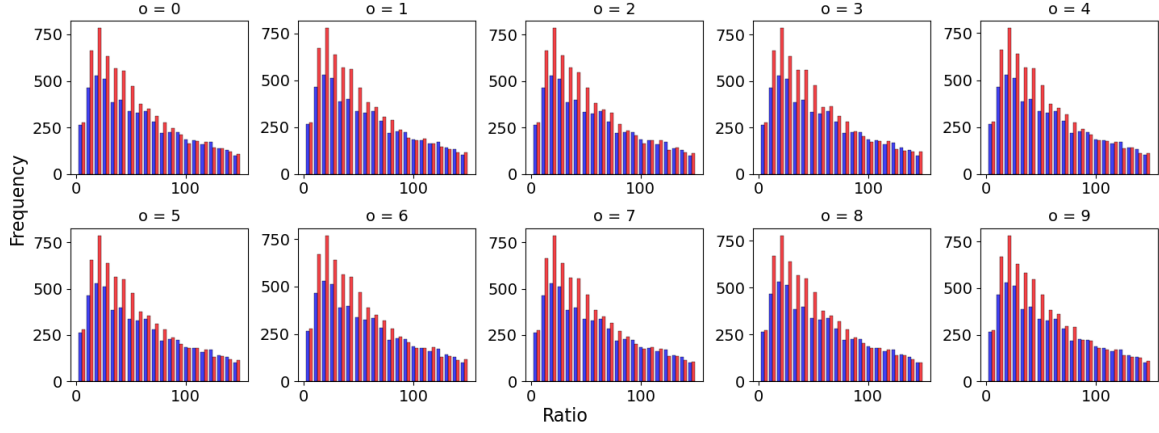


Figure 5: Histograms of $\tilde{r}_o(x)$ (blue bars) and $\tilde{r}'_o(x)$ (red bars) with range 1.4 to 150.

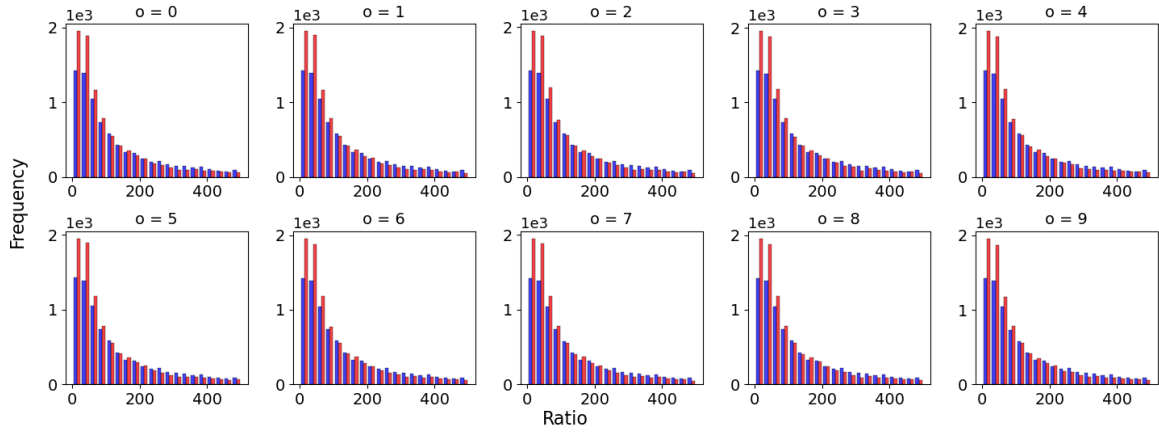


Figure 6: Histograms of $\tilde{r}_o(x)$ (blue bars) and $\tilde{r}'_o(x)$ (red bars) with range 1.4 to 500.

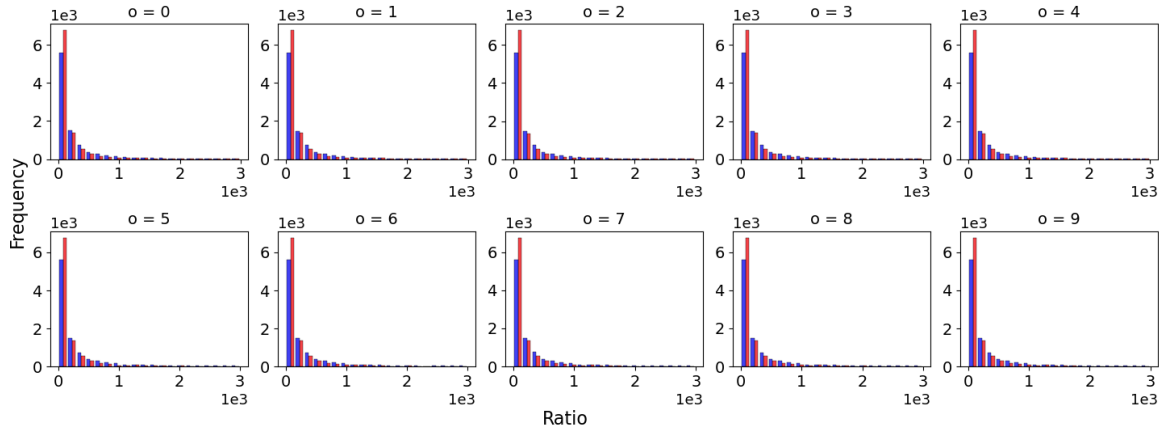


Figure 7: Histograms of $\tilde{r}_o(x)$ (blue bars) and $\tilde{r}'_o(x)$ (red bars) with range 1.4 to 3000.

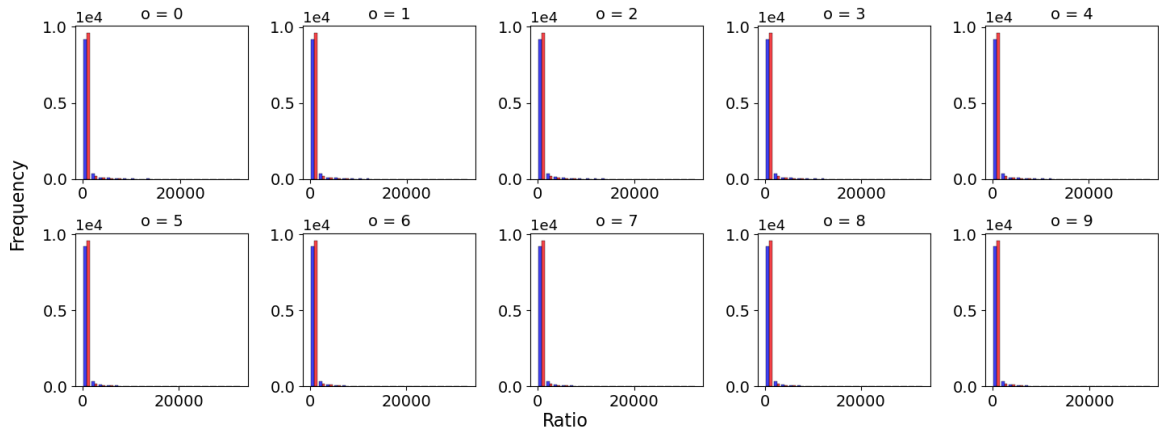


Figure 8: Full Histograms of $\tilde{r}_o(x)$ (blue bars) and $\tilde{r}'_o(x)$ (red bars) with range 1.4 to 33000.