

CROSS-ENTROPY LOSS OF APPROXIMATED DEEP NEURAL NETWORKS

Efficient Processing of Deep Neural Networks

(Sze, Chen, Yang, Emer, Morgan & Claypool Publishers, 2020)

The energy efficiency of DNNs on low-power, battery-operated embedded hardware (e.g., cellphones, smartwatches, augmented reality glasses) is a critical challenge.

→ **reducing the energy cost of DNNs:**

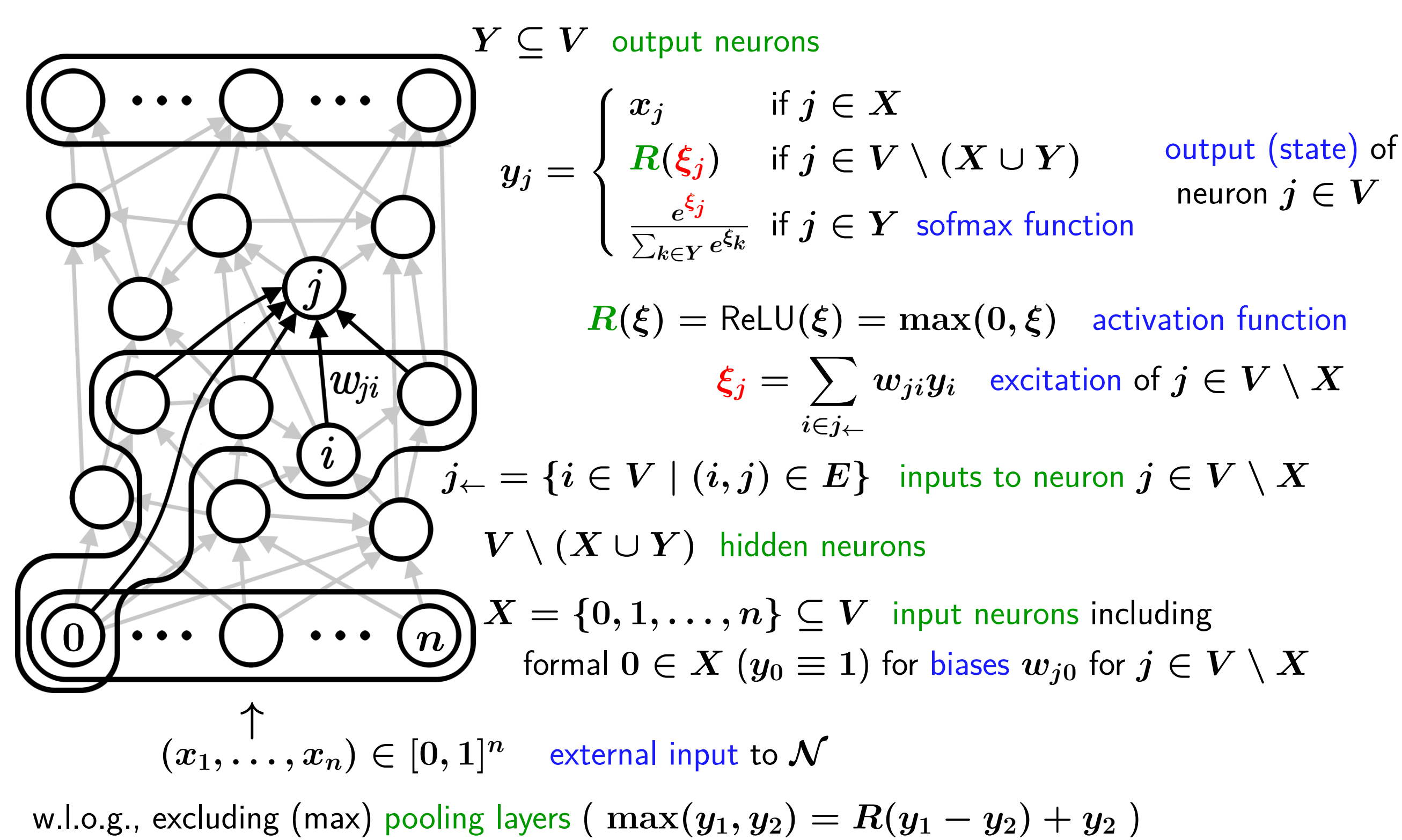
1. Specialized **Hardware Accelerators** for DNN inference (on GPUs, FPGAs, in-memory, etc.)
2. **Approximate Computing** in error-tolerant applications (e.g. image classification) save large amounts of energy with **minimal accuracy loss** by reducing:

- **Model Size:** pruning, compression, weight sharing, approximate multipliers
- **Arithmetic Precision:** fixed-point operations, reduced weight bit-width, nonuniform quantization

The aim of this study: Estimate the **maximum cross-entropy loss** of **approximated classification DNNs**.

Formal Model of DNN

The **architecture** of a DNN $\mathcal{N}$ is a connected directed acyclic graph (V, E) composed of neurons, where edges $(i, j) \in E \subset V \times V$ are labeled with real-valued **weights** w_{ji} :



1. Tool: Shortcut Weights

The excitation ξ_j of any neuron $j \in V \setminus X$ is a continuous **piecewise linear** function of the external input

$$\rightarrow \xi_j = \sum_{i \in X} w_{ji} y_i \text{ for } (y_1, \dots, y_n) \in \Xi$$

within a subset $\Xi \subset [0, 1]^n$ of the input space, where w_{ji} are referred to as the **shortcut weights** (bias) from **input neurons** $i \in X$ to neuron $j \in V \setminus X$.

For an input $(x_1, \dots, x_n) \in [0, 1]^n$, let $\Xi_S \subset [0, 1]^n$ be its **neighborhood** within which ξ_j are linear for all $j \in V \setminus X$ under **fixed shortcut weights**, where

$$S = S(x_1, \dots, x_n) = \{j \in V \setminus (X \cup Y) \mid \xi_j < 0\}$$

denotes the set of hidden neurons **saturated** at zero output, $y_j = R(\xi_j) = 0$.

→ Ξ_S is a **convex polytope**—an intersection of finitely many **half-spaces**:

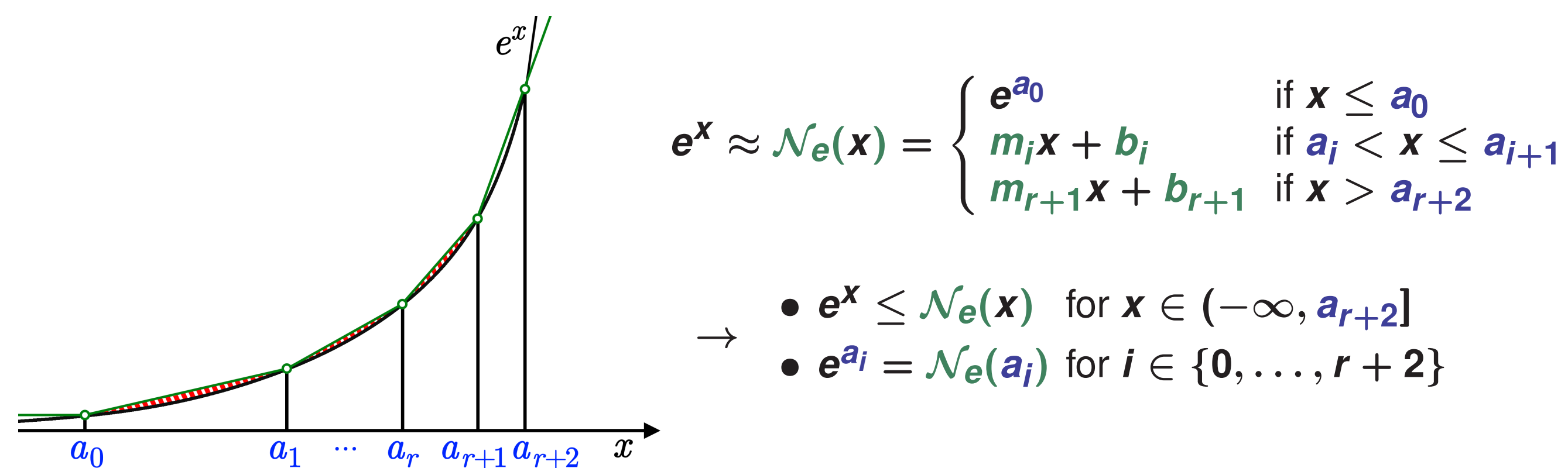
$$\xi_j = \sum_{i \in X} w_{ji} y_i \begin{cases} < 0 & \text{if } j \in S \\ \geq 0 & \text{if } j \notin S \end{cases} \text{ for } j \in V \setminus (X \cup Y)$$

$$0 \leq y_i \leq 1 \text{ for } i \in X$$

Efficient Computation of the shortcut weights via **feedforward propagation** (skipping S) for $j \in V \setminus X$:

$$w_{ji} = \sum_{k \in j^- \setminus S} w_{jk} w_{ki} \text{ for all } i \in X \quad \left(\text{initially } w_{ki} = \begin{cases} 1 & \text{if } k = i \\ 0 & \text{otherwise} \end{cases} \text{ for } k, i \in X \right)$$

2. Tool: Continuous Piecewise Linear Interpolation of e^x



Theorem: There are r unique points $a_1, \dots, a_r$ inside $[a_0, a_{r+1}]$ that minimize

$$\sum_{i=0}^r \int_{a_i}^{a_{i+1}} (m_i x + b_i - e^x) dx \quad \left(\rightarrow e^{a_i} = \frac{e^{a_{i+1}} - e^{a_{i-1}}}{a_{i+1} - a_{i-1}} \right).$$

Evaluating the **linear interpolation** using **ReLU networks**:

$$\mathcal{N}_e(x) = e^{a_0} + \sum_{i=0}^r R(m_i x + b_i - e^{a_i}) - \sum_{i=0}^r R(m_i x + b_i - e^{a_{i+1}})$$

Approximated DNNs

$\tilde{\mathcal{N}}$ is an **approximation** of $\mathcal{N}$, sharing the same input neurons ($\tilde{X} = X$) and the same number of output neurons ($|\tilde{Y}| = |Y|$) ($\tilde{\cdot}$ denotes parameters of $\tilde{\mathcal{N}}$),

e.g., the **weights** of $\tilde{\mathcal{N}}$ are **rounded** to a given **number of binary digits** in their **floating-point representations**

Cross-Entropy Loss of Approximated Classification DNNs

The **classification error** of $\tilde{\mathcal{N}}$ for an input $(x_1, \dots, x_n) \in [0, 1]^n$, measured by the **cross-entropy loss** between the **softmax** categorical probability distributions of $\mathcal{N}$ and $\tilde{\mathcal{N}}$:

$$L(x_1, \dots, x_n) = - \sum_{k \in Y} y_k \ln \tilde{y}_k \quad (\text{upper bounds the Kullback-Leibler divergence of } \mathcal{N} \text{ from } \tilde{\mathcal{N}})$$

Theorem: It is **NP-hard** to find the **maximum** of the **cross-entropy loss** L over the input space $[0, 1]^n$.

Upper-Bounding L Using ReLU Networks $\tilde{\mathcal{N}}_{ck}$ & $\mathcal{N}_c$

within the **convex polytope** Ξ_S around a **data point** $(x_1, \dots, x_n) \in T$ of **category** $c \in Y$ from a **test/training set** $T \subset [0, 1]^n$, where $S = S(x_1, \dots, x_n)$, restricted to inputs in Ξ_S that are classified by $\mathcal{N}$ into the category c with probability at least p (e.g., $p = 0.8$): $y_c \geq p$

$$\sum_{j \in Y \setminus \{c\}} e^{\xi_j - \xi_c} \leq \frac{1 - y_c}{y_c} \leq \frac{1}{p} - 1 \quad \xrightarrow{\text{softmax}} y_c \geq p$$

$$L(x_1, \dots, x_n) = \sum_{k \in Y} y_k \ln \frac{1}{y_k} \leq y_c \ln \frac{1}{y_c} + (1 - y_c) \max_{k \in Y \setminus \{c\}} \ln \frac{1}{y_k}$$

$$\stackrel{\text{softmax}}{=} \max_{k \in Y \setminus \{c\}} \ln \sum_{j \in Y} e^{\tilde{\xi}_j - \tilde{\xi}_k + y_c (\tilde{\xi}_k - \tilde{\xi}_c)} \stackrel{y_c \geq p}{\leq} \ln \max_{k \in Y \setminus \{c\}} \sum_{j \in Y} e^{\tilde{\xi}_j - \tilde{\xi}_k + R(\tilde{\xi}_k - \tilde{\xi}_c) - pR(\tilde{\xi}_c - \tilde{\xi}_k)}$$

$$\stackrel{(*)}{\leq} \ln \max_{k \in Y \setminus \{c\}} \sum_{j \in Y} \mathcal{N}_e(\tilde{\xi}_j - \tilde{\xi}_k + R(\tilde{\xi}_k - \tilde{\xi}_c) - pR(\tilde{\xi}_c - \tilde{\xi}_k)) = \ln \max_{k \in Y \setminus \{c\}} \tilde{\mathcal{N}}_{ck}(x_1, \dots, x_n)$$

AppMax Method

Input: DNN $\mathcal{N}$, its approximation $\tilde{\mathcal{N}}$, $(x_1, \dots, x_n) \in T$ of category $c \in Y$

Output: an **upper bound** on

$$\max_{(y_1, \dots, y_n) \in \bigcup_{k \in Y \setminus \{c\}} \Xi_k^*} L(y_1, \dots, y_n)$$

Algorithm: For each $k \in \tilde{Y} \setminus \{c\}$ do

- Compose $\tilde{\mathcal{N}}_{ck}^*$ from $\tilde{\mathcal{N}}_{ck}$ & $\mathcal{N}_c$.
- Determine the **saturated neurons** $S^* = S^*(x_1, \dots, x_n)$ in $\tilde{\mathcal{N}}^*$.
- Compute the **shortcut weights** w_{ji}^* of $\tilde{\mathcal{N}}^*$ for all $j \in V^* \setminus X^*$ and $i \in X^*$.
- Solve the **linear program** (LP) to find $(y_1, \dots, y_n)$ that

$$\text{maximize } \tilde{\mathcal{N}}_{ck}(y_1, \dots, y_n) \quad (\rightarrow \sigma_k) \quad \text{over the polytope } \Xi_k^* \subseteq \Xi_S^*$$

defined by $\xi_j^* = \sum_{i \in X} w_{ji}^* y_i \begin{cases} \leq 0 & \text{if } j \in S^* \\ \geq 0 & \text{if } j \notin S^* \end{cases} \text{ for } j \in V^* \setminus (X^* \cup Y^*),$

$$\tilde{\mathcal{N}}_c(x_1, \dots, x_n) \leq \frac{1}{p} - 1 \quad (\text{where } e^{a_0} \leq \frac{1}{p} - 1 \leq e^{a_{r+2}}),$$

$$\tilde{\xi}_j - \tilde{\xi}_k + R(\tilde{\xi}_k - \tilde{\xi}_c) - pR(\tilde{\xi}_c - \tilde{\xi}_k) \leq a_{r+2} \quad (*), \quad (y_1, \dots, y_n) \in [0, 1]^n.$$

and output $\ln \max_{k \in \tilde{Y} \setminus \{c\}} \sigma_k$

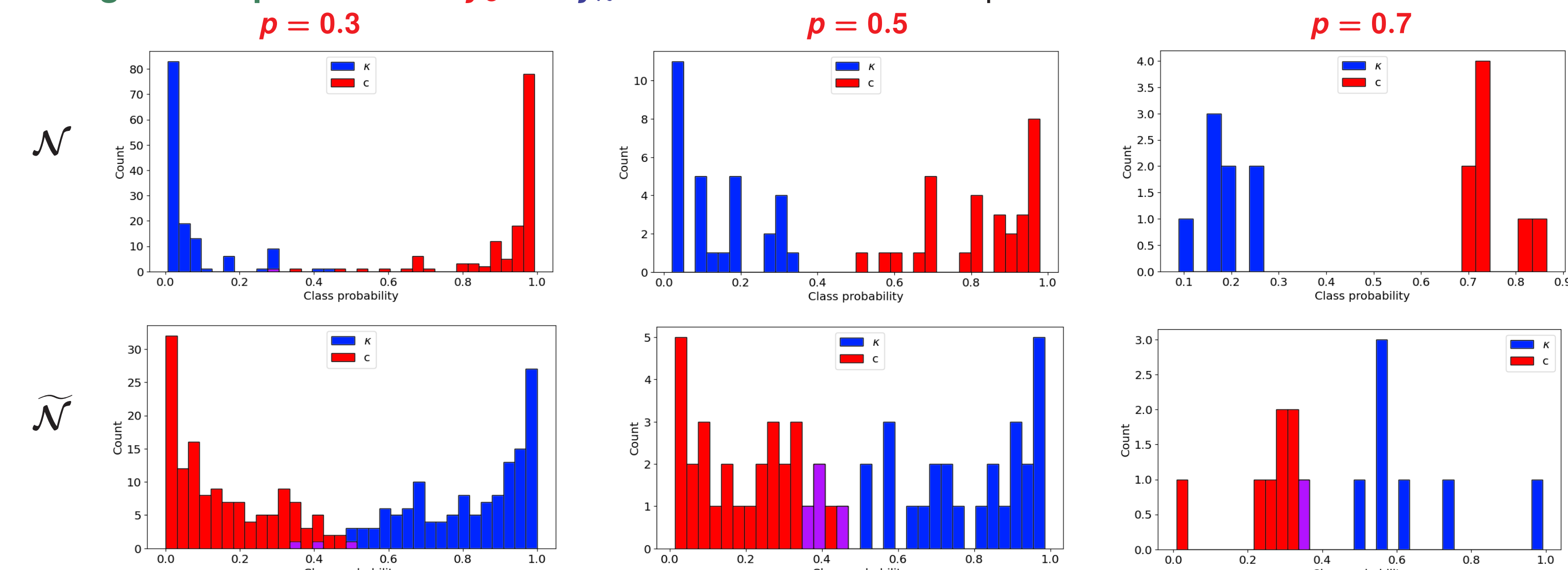
Experiments

- **DNN:** $\mathcal{N}$ with 3 fully connected layers (784–64–32–10) and 32-bit weights, trained on the **MNIST** dataset; approximated as $\tilde{\mathcal{N}}$ with weights rounded to 4 bits
- **Source Code:** <https://github.com/PetraVidnerova/ClassificationRoundingErrors>
- **Test Set:** T with 1135 data points of class $c = 1$, on which $\mathcal{N}$ and $\tilde{\mathcal{N}}$ achieve 100% and 99.91% **accuracy**, respectively
- **Linear Interpolation of e^x :** $r = 14$ optimal points & $a_0 = -5$, $a_{r+1} = 5$, $a_{r+2} = 20$

Number of “**misclassified polytopes**” around data points in T , including so-called **misclassified inputs** with the **maximum estimated upper bound of cross-entropy loss**, which are classified correctly as $c = 1$ by $\mathcal{N}$ with **probability at least p** , but misclassified as $\kappa \neq c$ by $\tilde{\mathcal{N}}$:

p	0.30	0.40	0.50	0.60	0.70	0.80	0.90	0.95	0.99
misclassified	134	71	30	11	8	0	0	0	0

Histograms of probabilities y_c and y_{κ} over the misclassified inputs:



Future Research Directions

- Broaden AppMax evaluation to other **datasets** (e.g., CIFAR-100, ImageNet).
- Most suitable error metric (KL **divergence**, **cross-entropy loss**, **accuracy**) with optimal upper bound.
- **AppMax** for **classification DNNs** with **softmax**, via nonlinear **Karush-Kuhn-Tucker** optimization.
- Approximate **global error** by estimating the **probabilities of convex polytopes** from their **volumes**, measured using the average **mean width** (evaluated by LP).
- Identify **DNN components that can be neglected** (e.g., specific weights to be rounded) under explicitly bounded output error.