
Hard Samples, Bad Labels: Robust Loss Functions That Know When to Back Off

Nicholas Pellegrino^{1,*}, David Szczecina^{1,2,*}, & Paul Fieguth¹

¹Vision and Image Processing Group, Systems Design Engineering, University of Waterloo

²Mechanical & Mechatronics Engineering, University of Waterloo

{npellegr,dszczeci,pfieguth}@uwaterloo.ca

Abstract

Incorrectly labelled training data are frustratingly ubiquitous in both benchmark and specially curated datasets. Such mislabelling clearly adversely affects the performance and generalizability of models trained through supervised learning on the associated datasets. Frameworks for detecting label errors typically require well-trained / well-generalized models; however, at the same time most frameworks rely on training these models on corrupt data, which clearly has the effect of reducing model generalizability and subsequent effectiveness in error detection — unless a training scheme robust to label errors is employed. We evaluate two novel loss functions, *Blurry Loss* and *Piecewise-zero Loss*, that enhance robustness to label errors by de-weighting or disregarding difficult-to-classify samples, which are likely to be erroneous. These loss functions leverage the idea that mislabelled examples are typically more difficult to classify and should contribute less to the learning signal. Comprehensive experiments on a variety of artificially corrupted datasets demonstrate that the proposed loss functions outperform state-of-the-art robust loss functions in nearly all cases, achieving superior F1 scores for error detection. Further analyses through ablation studies offer insights to confirm these loss functions’ broad applicability to cases of both uniform and non-uniform corruption, and with different label error detection frameworks. By using these robust loss functions, machine learning practitioners can more effectively identify, prune, or correct errors in their training data. Code, including a working demonstration Jupyter Notebook, is available at https://anonymous.4open.science/r/Robust_Loss-6BAD/.

1 Introduction

Supervised learning relies heavily on the assumption that training labels are accurate; however, label errors are pervasive even in well-established benchmark datasets, typically at rates of at least 5% [26], including in classic datasets such as MNIST [4, 26]. Label errors (also referred to as label *noise* in the literature) degrade model performance, limit generalizability [31, 33, 42], and mislead validation metrics. The prevalence of label errors is especially problematic in domains with hierarchical and fine-grained classes, such as biological datasets [5–7, 12, 24, 34, 39], where mislabelling often occurs among highly similar categories. Efficient automatic error detection could significantly streamline curation and the flagging of samples for expert review, reducing costs and improving dataset quality.

Existing work on detecting label errors [13, 16, 25, 31, 32] such as Confident Learning (CL) [25] and Area Under Margin (AUM) [31] typically rely on training surrogate models, intended to be robust or resistant to overfitting, such that they are well-generalized and not fit to erroneous labels in the training data. Crucially, the effectiveness of these methods depends on the models producing statistically distinguishable predicted probabilities, $p(k|x)$, for erroneously *vs.* correctly labelled

*Indicates equal contribution, joint first-authorship.

samples. However, when trained with standard loss functions such as Cross Entropy [10] or Focal Loss [19], models tend to inadvertently fit to erroneous labels, reducing detection effectiveness. Modification of the loss function used when training such models is one avenue for improving their robustness and downstream utility for detecting label errors.

In this paper, we rigorously investigate two recently introduced robust loss functions — *Blurry Loss* and *Piecewise-zero Loss* [30] — originally presented with only very few and preliminary studies performed, and specifically designed to improve model robustness by explicitly de-weighting samples likely to be mislabelled. In datasets with label errors, erroneous samples are likely to be difficult-to-classify. Inspired by Focal loss [19], which places emphasis on difficult-to-classify samples, these novel loss functions reverse this idea, assigning less weight to samples with low predicted probability in their as-labelled class (likely to be mislabelled), thereby avoiding detrimental overfitting to erroneous data. These loss functions result in a more robust training scheme and models that are better suited to detecting label errors than those trained with existing standard or robust loss functions.

Comprehensive experiments are performed to validate these loss functions’ effectiveness in improving label error detection, exploring several artificially corrupted datasets, a range of artificial corruption rates, uniform and non-uniform error statistics, multiple frameworks for label error detection (CL and AUM), and gradients of corrupted and clean training samples. The proposed Blurry and Piecewise-zero loss functions outperform existing loss functions, Generalized Cross Entropy [43] (GCE) and Active-Negative Loss [41] (ANL), across nearly all experiments, highlighting their utility for practical error detection and correction.

2 Preliminaries

Building upon [41, 43], this work is set in the context of a K -class classification problem based with $\mathcal{X} \subset \mathbb{R}^r$ as the r -dimensional feature space (typically images) and $\mathcal{Y} = \{1, \dots, K\}$ as the label space (categorical). A clean (ideal, error-free) dataset is denoted by $D = \{(x_i, y_i)\}_{i=1}^N$ for sets containing N examples and where $(x_i, y_i) \in (\mathcal{X}, \mathcal{Y})$. A classifier is a function, $f : \mathcal{X} \rightarrow \mathcal{Y}$, that maps input features to predicted classes by selecting the class, k , having the maximal predicted probability, $f = \arg_k \max p(k|x)$. The predicted probabilities, $p(k|x)$, for each class $k \in \mathcal{Y}$, are inferred through a trainable deep neural network (DNN) with a softmax output layer, such that $\sum_{k=1}^K p(k|x) = 1$. In practice, classifiers are trained through the strategy of attempting to minimize the empirical risk, $R_{\mathcal{L}} = \frac{1}{N} \sum_{i=1}^N \mathcal{L}(p(k|x_i), y_i)$ for a given loss function $\mathcal{L} : [0, 1]^K \times \mathcal{Y} \rightarrow \mathbb{R}^+$.

Label errors are often present within training data, as discussed in Section 1. The proportion of data that are mislabelled is termed the corruption rate, η , and the corrupted dataset is denoted $D_{\eta} = \{(x_i, \tilde{y}_i)\}_{i=1}^N$ for possibly erroneous labels $\tilde{y}_i$. The probability that a correct label y is mislabelled (corrupted) as label $\hat{y}$ is $\eta_{y\hat{y}}$, such that the corruption rate of specific label y is given by $\eta_y = \sum_{\hat{y} \neq y} \eta_{y\hat{y}}$. Taking the common assumption that the label corruption (*i.e.*, mislabelling) process is conditionally independent of input features [23], x , labels are given by

$$\tilde{y} = \begin{cases} y & \text{with probability } (1 - \eta_y) \\ \hat{y}, \hat{y} \in \mathcal{Y} \setminus \{y\}, & \text{with probability } \eta_{y\hat{y}} \end{cases}. \quad (1)$$

Corruption processes may be categorized based on their conditioning. In the most general case of *asymmetric* corruption, the corruption process is class-dependent and conditioned on both correct labels, y , and corrupted labels, $\hat{y}$. Corruption is *symmetric* if the process is *identically* conditioned on both correct and incorrect labels. Lastly, the corruption process is *uniform* if it is not conditioned on either correct or incorrect labels.

For the sake of notational simplicity throughout the remainder of this article, predicted probabilities, $p(k|x)$, are stated in terms of an *implied* input x , such that the notational simplification is made, from $p(k|x)$ to p_k (indicating predicted probability for class k).

3 Background

The background information is divided into two main categories: Label Error Detection, and Robust Loss Functions. Because the label error detection frameworks of CL [25] and AUM [31] are used in this paper to evaluate the proposed loss functions, more detail is given to them than to other

frameworks. Similarly, the robust loss functions Generalized Cross Entropy (GCE) [43] and Active Negative Loss (ANL) [41] are used in this paper and presented in more detail.

3.1 Label Error Detection

The prevalence and negative impacts on training models caused by errors in training data have motivated the development of label error detection frameworks. Two prominent approaches in this space are Confident Learning (CL) [25] and Area Under the Margin (AUM) [31], both of which rely on training surrogate models on corrupt datasets in order to detect samples likely to have label errors. Other recent frameworks include MentorNet [13], SELFIE [32], and SelNLPL (Selective Negative Learning and Positive Learning) [16].

Confident Learning [25] uses k-fold cross-validation to train models on all folds of the mislabelled dataset, wherein, for each fold, the joint distribution of observed, $\tilde{y}$, and underlying true labels, y , termed the Confident joint, $Q_{\tilde{y}, y}$, is estimated based on the model’s predicted probabilities. Proposal methods for sample pruning available in the framework include Prune by Class (PBC), Prune by Noise Rate (PBNR), and “both”. PBC detects examples from each class that the model is least confident about, whereas PBNR detects examples that are most likely to be mislabelled as a *different* class. With PBC, for each class k , examples with the lowest self-confidence (*i.e.*, $p(k|x)$), for inputs x labelled class k) are proposed for pruning. With PBNR, examples from each class k having the largest margin, (*i.e.*, $p(j|x) - p(k|x)$, $j \neq k$), are proposed for pruning. If the “both” method is selected, examples must both be proposed by PBC and PBNR to be detected.

The Area Under Margin [31] method identifies mislabelled samples based on their training dynamics. AUM measures the difference between the logit values for a training example’s as-labelled class, $\tilde{y}$, and its highest other class, and averages over the course of training. Correctly labelled samples tend to exhibit positive AUM values, while mislabelled samples tend to have significantly lower AUM values due to conflicting gradient updates from other samples of the same class. To determine at which threshold of the AUM statistic samples should be detected as being likely mislabelled, a small portion of the data are re-labelling to a new *fake* class, numbered $K + 1$ for a K -class problem, and correspondingly one additional output neuron is added to the DNN. The AUM for this set of ‘threshold samples’ is monitored, and the 99th percentile AUM from this set is used as the threshold at which correctly- and mislabelled data are separated.

3.2 Robust Loss Functions

The choice of loss function largely determines the training dynamics and the generalization properties of models trained under supervision. For classification problems, Categorical Cross Entropy (CE) [10] loss is standard. For predicted probabilities p_k , as-labelled class $\tilde{y}$, Cross Entropy loss is defined as $CE(p_k, \tilde{y}) = -\log(p_{\tilde{y}})$. Focal loss (FL) [19] builds upon Cross Entropy loss by placing additional weight on difficult-to-classify samples, through the inclusion of an additional factor, $(1 - p_{\tilde{y}})^\gamma$, with weighting parameter γ (with a standard setting of $\gamma = 2$). Focal loss is defined as $FL(p_k, \tilde{y}) = -(1 - p_{\tilde{y}})^\gamma \log(p_{\tilde{y}})$. Note that if $\gamma = 0$, Focal loss reduces simply to Cross Entropy. Despite their widespread use, standard loss functions are sensitive to label errors [9, 28, 33] as a result of the strong penalization of confident mispredictions, wherein $p_{\tilde{y}}$ is low for erroneous $\tilde{y} = \hat{y}$.

Robustness to label errors has motivated several alternative loss functions. Mean Absolute Error (MAE) is theoretically robust to symmetric label corruption [8, 9], but has been shown to perform poorly on complicated datasets [43]. Generalized Cross Entropy (GCE) [43] combines Cross Entropy with MAE to obtain the positive characteristics of both losses. The GCE loss function, officially termed $\mathcal{L}_q$ loss, is defined as the negative Box-Cox transformation[1],

$$\text{GCE}(p_k, \tilde{y}) = \mathcal{L}_q(p_k, \tilde{y}) = \frac{(1 - p_{\tilde{y}}^q)}{q}, \quad (2)$$

where $q \in (0, 1]$. The parameter q controls the transition between these two losses, such that in the limit as $q \rightarrow 0$, GCE becomes Cross Entropy (via L’Hôpital’s rule), and for $q = 1$, GCE becomes MAE. The recommended parameter setting is $q = 0.7$ for most problems [43].

Active Negative Losses (ANL) [41] set the current state-of-the-art and are a class of loss functions that combine Normalized Loss Functions [21], $\mathcal{L}_{\text{norm}}$, (originally introduced to be robust to label errors;

normalizes by dividing loss for as-labelled class with loss summed over all classes) and Normalized Negative Loss Functions (NNLFs), $\mathcal{L}_{nn}$, (introduced in [41]) through weighting parameters, $a, b > 0$,

$$\text{ANL}(p_k, \tilde{y}) = \alpha \cdot \mathcal{L}_{\text{norm}} + \beta \cdot \mathcal{L}_{nn}. \quad (3)$$

Any existing loss function, such as Cross Entropy or Focal loss, can be converted to an ANL through the substitution of normalized and normalized negative forms. The inclusion of L1 regularization is suggested to avoid overfitting, which commonly occurs for ANL without regularization.

Other notable robust loss functions include Symmetric Cross Entropy [37] (combines Reverse Cross Entropy [27] with CE), PHuber-CE [22] (composite loss-based gradient clipping applied to CE), Active Passive Loss (APL) [21] (designed to maximize $p_{\tilde{y}}$ while minimizing $p_k, k \neq \tilde{y}$), Asymmetric Loss Functions (ALFs) [44], and Curriculum Loss [20]. These precursors set the stage for continued research into robust loss design.

4 Method

Two novel loss functions are introduced that are designed to be robust to label errors. During training on corrupt data, $D_\eta = \{(x_i, \hat{y}_i)\}_{i=1}^N$, mislabelled samples $(x, \hat{y})$ are likely difficult to classify and thus have low predicted probabilities, $p_{\tilde{y}} = p_{\hat{y}}$, for their as-labelled class, $\hat{y}$. With CE loss, having low predicted probability, $p_{\tilde{y}}$, means the *gradient* of the loss is *large and negative*, imparting a strong signal to the optimizer causing fitting to these samples. The proposed loss functions achieve their robustness by reducing the loss for samples likely to be mislabelled, either resulting in a *positive* or *zero* gradient and no longer incorrectly steering the optimizer towards fitting to these erroneous samples.

Though the proposed loss functions are the underlying usable contributions to the field, the greater contribution of this article is in the thorough comparative evaluation against other state-of-the-art losses and ablation studies used to provide a deeper understanding of their mechanisms. These experiments involve the use of these loss functions within existing label error detection frameworks, Confident Learning (CL) and Area Under the Margin (introduced in Section 3.1), to correctly identify mislabelled data in artificially corrupted datasets, and compare the effectiveness resulting from the use of these proposed loss functions *vs.* existing state-of-the-art robust loss functions. The proposed loss functions are outlined here, whereas the experimental details are outlined in Section 5.

4.1 Blurry Loss

Blurry Loss [30] is closely related to Focal loss [19], but *de-weights* difficult-to-classify samples, likely to be mislabelled, through a multiplicative factor. For predicted probability p_k , as-labelled class $\tilde{y}$, and a weighting parameter $\gamma \in \mathbb{R}$, Blurry Loss is defined as

$$\text{BL}(p_k, \tilde{y}) = -(p_{\tilde{y}})^\gamma \log(p_{\tilde{y}}). \quad (4)$$

Figure 1a shows the this loss function for several variations of parameter γ .

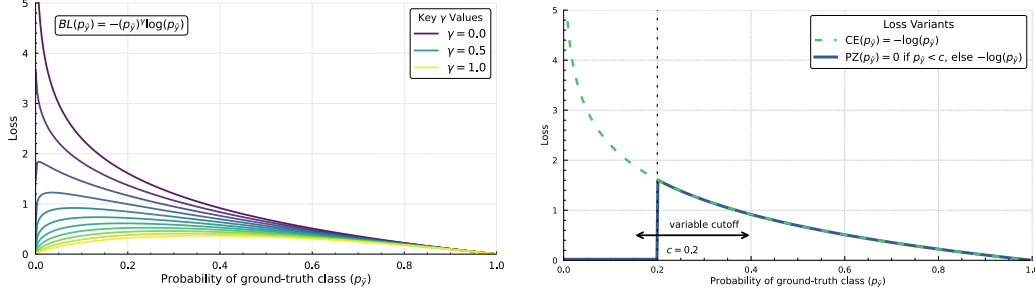
Notice that this function is non-monotonic and features a section of *positive* gradient for $p_{\tilde{y}} < e^{-1/\gamma}$, encouraging training *against* the as-labelled class if $p_{\tilde{y}}$ is low. At $\gamma = 0$, Blurry Loss is equivalent to Cross Entropy Loss.

4.2 Piecewise-zero Loss

Piecewise-zero Loss [30] is a variant of CE loss and is designed to ignore (zero) difficult-to-classify samples, which are likely to be mislabelled. Examples with predicted probability beneath some cutoff, $p_{\tilde{y}} \leq c \in [0, 1]$, are assigned a loss of zero through a piecewise definition of the function. Also note that the *gradient* within the cutoff region is also zero, causing these examples with low $p_{\tilde{y}}$ not to affect the training process (weights are not updated by zero-gradient samples). Piecewise-zero Loss is defined as

$$\text{PZ}(p_k, \tilde{y}) = \begin{cases} 0 & p_{\tilde{y}} \leq c, \\ \text{CE}(p_k, \tilde{y}) = -\log(p_{\tilde{y}}) & p_t > c. \end{cases} \quad (5)$$

Figure 1b shows this loss function for several variations of the cutoff-position parameter c . At $c = 0$, Piecewise-zero Loss is equivalent to Cross Entropy Loss.



(a) The Blurry Loss function of Equation (4), plotted for a range of parameter γ , controlling the shape. Note that at $\gamma = 0$, this is equivalent to CE. (b) The Piecewise-zero Loss function of Equation (5), illustrating the effect of adjusting the cutoff position, c . Note that at $c = 0$, this is equivalent to CE.

Figure 1: The two proposed loss functions: Blurry Loss (a) and Piecewise-zero Loss (b).

4.3 Loss Scheduling

DNN model parameters are typically randomly initialized at the start of training. As a result, the model’s outputs at the early stages of training bear no meaningful relationship to the training labels. Applying the proposed loss functions immediately — especially the Piecewise-zero Loss with a high cutoff — is likely to lead to many correctly labelled samples being de-weighted or ignored, dramatically reducing the ability to train. Instead, a scheme through which training begins with Cross Entropy loss is proposed, to allow the model to sufficiently generalize prior to the use of these novel loss functions. A delay parameter, $d \in \mathbb{N}$, determines the epoch at which the training switches to one of the proposed loss functions. With this scheduling approach, for example with PZ loss, the loss, $\mathcal{L}_d$, used is given by

$$\mathcal{L}_d(p_k, \tilde{y}) = \begin{cases} \text{CE}(p_k, \tilde{y}) & \text{Epoch} \leq d, \\ \text{PZ}(p_k, \tilde{y}) & \text{Epoch} > d. \end{cases} \quad (6)$$

5 Experiments

The following paragraphs outline the datasets, method of artificial corruption, loss functions used for comparison, models and training scheme, and metrics used for the experiments.

Datasets and Artificial Corruption Synthetic uniform label noise is injected into standard benchmark datasets following the same method used in [23, 29, 41, 43]. For a given corruption rate $\eta \in [0, 1]$, a fraction η of the training labels are randomly flipped. Corruption rates in $\eta \in \{0.1, 0.2, 0.3, 0.4\}$ are tested. Datasets include MNIST [4], Fashion-MNIST [40], CIFAR-10, and CIFAR-100 [18], under the assumption that their existing label error rates are negligible compared to the induced corruption, η .

To assess robustness under realistic, non-uniform noise, the CIFAR-10N and CIFAR-100N [38] datasets are employed, which are versions of the CIFAR-10 and CIFAR-100 datasets, re-annotated using Amazon Mechanical Turk. By design, the re-annotations do *not* correct labels but rather induce abundant, real-world, non-uniform human annotation errors. The original CIFAR labels are considered ground truth in the CIFAR-N datasets. For CIFAR-10N, three re-annotations are given, whereas for CIFAR-100, only one is given (with $\eta = 0.4020$). With CIFAR-10N, we use the ‘Aggregate’ version (labels used are the majority vote of the three re-annotations; $\eta = 0.0903$) and the ‘Worst’ version (re-annotations *not* matching ground truth are always used when available; $\eta = 0.4021$).

Baseline Loss Functions Categorical Cross-Entropy [10] and Focal Loss [19] are the primary baseline loss functions, as these serve as the *de facto* gold standard for most supervised classification problems. Focal Loss is used with $\gamma = 2$ as recommended in [19]. To assess robustness under label errors, we also include state-of-the-art robust loss functions, Generalized Cross-Entropy [43] (GCE)

loss, defined in Equation (2), Active Negative Cross-Entropy [41] (ANL-CE) and Active Negative Focal Loss [41] (ANL-FL), defined in Equation (3).

Consistent with the original formulation, the GCE hyperparameter q is fixed at 0.7 as is recommended in [43]. For the ANL variants, the dataset-specific hyperparameters α , β , and δ , as well as the focal-loss scaling factor γ for ANL-FL, are set exactly as prescribed in [41]. However, since the Fashion-MNIST dataset was not used in [41], no parameter specification was provided and thus the MNIST parameters of [41] are used for the Fashion-MNIST experiments here. Each loss function is evaluated using these recommended settings to ensure a fair comparison.

Models and Hyperparameters For MNIST and Fashion-MNIST, we adopt a minimal convolutional neural network consisting of two convolutional layers followed by two fully connected layers interleaved with dropout. Under a ten-epoch training schedule, this architecture reliably exceeds 99% test accuracy on clean data (near state-of-the-art) [14]. For CIFAR-10, CIFAR-100, and their noisy variants (CIFAR-10N/CIFAR-100N), we fine-tune an ImageNet-pretrained ResNet-34 [11] gently modified for lower resolution images. After training for ten epochs, the model attains 93.5% test accuracy on CIFAR-10 and 74.5% on CIFAR-100. Additional details on models and optimizers are given in Section A.1.

To explore the parameter-sensitivity of our proposed losses, we perform grid sweeps over the Blurry Loss exponent, γ , and the Piecewise-Zero Loss cutoff, c , thereby characterizing their influence on label error detection efficacy and allowing optimal values to be identified for each problem. As demonstrated in ablation Figure 5 of Section A.2, introducing a one-epoch delay yields the best performance for PZ loss; accordingly, we use $d = 1$ for all subsequent PZ loss experiments.

Label Error Detection and Performance Metrics To evaluate each loss function’s efficacy for the detection of label errors, we apply the Confident Learning (CL) [25] and AUM [31] frameworks, both described in Section 3.1. The CL pruning method used in all cases is “both”, in which detected samples must both be unlikely to be their as-labelled class and likely to be of another different class. Five cross-validation folds are used with 80% train / 20% predict splits. The use of AUM serves as an ablation study of label error detection, to demonstrate the efficacy of the proposed loss functions across multiple detection frameworks. Following the methodology of [31], experiments on CIFAR-100 are replicated, using a ResNet-32 [11] model trained for 150 epochs with a learning rate of 0.1 and weight decay of 10^{-4} ; however, with Cross-Entropy (CE) [10] loss replaced with Blurry and Piecewise Zero Loss for a direct comparison.

Performance is primarily quantified using the F1 score [35] — a gold-standard metric defined as the harmonic mean of precision and recall. The F1 score penalizes situations in which either precision or recall fall, which is not always reflective of one’s goals when attempting to identify label errors. To address this limitation with low corruption rates, we also report the Balanced Accuracy score [2], defined as the average of sensitivity (true positive rate) and specificity (true negative rate), which provides a more robust measure of detector performance under class imbalance.

Compute Resource Requirements All experiments were conducted on a single NVIDIA Tesla V100-SXM2 GPU (16 GB VRAM). Reproducing the main CL label error detection results requires about 6 GB of GPU memory and takes roughly 10 minutes per loss-function and hyperparameter combination. AUM experiments demand 7 GB of memory and approximately 2.5 hours per loss-function and hyperparameter test. Further details are given in Section A.1.

6 Results

Results for the main experiment (Section 6.1) and additional studies (Section 6.2) are presented here. All tables are formatted such that best scores are **bolded and underlined**, and second best are simply **bolded**.

6.1 Main Findings

The proposed loss functions are compared against existing standard and robust loss functions, including CE, FL, GCE, ANL-CE, and ANL-PZ. This comparison is performed on several datasets and corruption rates, as described in Section 5. Results from 20 random trials are summarized in

Table 1: Label error detections, averaged over 20 random trials, measured using the F1 score. The parameters (γ and c) resulting in the best performance for each dataset and corruption rate are given for the proposed loss functions, Blurry Loss (BL) and Piecewise-zero Loss (PZ).

Dataset	η	CE	FL	GCE	ANL-CE	ANL-FL	BL	γ	PZ	c
MNIST	0.10	0.969 $\pm$ 0.002	0.960 $\pm$ 0.004	0.957 $\pm$ 0.002	0.941 $\pm$ 0.002	0.940 $\pm$ 0.002	0.979$\pm$0.002	0.3	0.974$\pm$0.001	0.005
	0.20	0.975 $\pm$ 0.001	0.967 $\pm$ 0.002	0.975 $\pm$ 0.001	0.970 $\pm$ 0.001	0.969 $\pm$ 0.001	0.984$\pm$0.001	0.4	0.983$\pm$0.001	0.010
	0.30	0.975 $\pm$ 0.001	0.967 $\pm$ 0.001	0.983 $\pm$ 0.001	0.980 $\pm$ 0.001	0.980 $\pm$ 0.001	0.988$\pm$0.001	0.4	0.986$\pm$0.001	0.020
	0.40	0.973 $\pm$ 0.001	0.963 $\pm$ 0.002	0.987 $\pm$ 0.001	0.985 $\pm$ 0.000	0.985 $\pm$ 0.001	0.990$\pm$0.000	0.5	0.988$\pm$0.000	0.040
Fashion-MNIST	0.10	0.824 $\pm$ 0.003	0.831$\pm$0.004	0.737 $\pm$ 0.004	0.683 $\pm$ 0.003	0.679 $\pm$ 0.004	0.826$\pm$0.003	0.1	0.818 $\pm$ 0.003	0.005
	0.20	0.884 $\pm$ 0.002	0.879 $\pm$ 0.003	0.850 $\pm$ 0.002	0.822 $\pm$ 0.002	0.819 $\pm$ 0.002	0.887$\pm$0.002	0.2	0.885$\pm$0.003	0.005
	0.30	0.907 $\pm$ 0.002	0.899 $\pm$ 0.001	0.900 $\pm$ 0.001	0.881 $\pm$ 0.002	0.879 $\pm$ 0.001	0.917$\pm$0.001	0.4	0.916$\pm$0.001	0.020
	0.40	0.918 $\pm$ 0.002	0.896 $\pm$ 0.033	0.927 $\pm$ 0.001	0.915 $\pm$ 0.001	0.914 $\pm$ 0.002	0.934$\pm$0.001	0.5	0.933$\pm$0.001	0.040
CIFAR-10	0.10	0.740 $\pm$ 0.006	0.757 $\pm$ 0.006	0.773 $\pm$ 0.006	0.698 $\pm$ 0.006	0.693 $\pm$ 0.007	0.787$\pm$0.005	0.5	0.794$\pm$0.008	0.040
	0.20	0.773 $\pm$ 0.004	0.789 $\pm$ 0.004	0.858 $\pm$ 0.004	0.821 $\pm$ 0.004	0.817 $\pm$ 0.005	0.860$\pm$0.004	0.6	0.863$\pm$0.004	0.060
	0.30	0.771 $\pm$ 0.004	0.798 $\pm$ 0.003	0.893 $\pm$ 0.003	0.873 $\pm$ 0.003	0.871 $\pm$ 0.003	0.894$\pm$0.003	0.7	0.894$\pm$0.003	0.080
	0.40	0.769 $\pm$ 0.003	0.801 $\pm$ 0.003	0.909 $\pm$ 0.003	0.901 $\pm$ 0.002	0.899 $\pm$ 0.003	0.912$\pm$0.003	0.8	0.912$\pm$0.003	0.100
CIFAR-100	0.10	0.503 $\pm$ 0.009	0.520$\pm$0.009	0.485 $\pm$ 0.009	0.399 $\pm$ 0.007	0.397 $\pm$ 0.007	0.507 $\pm$ 0.009	0.3	0.510$\pm$0.009	0.010
	0.20	0.628 $\pm$ 0.006	0.638 $\pm$ 0.006	0.652 $\pm$ 0.007	0.578 $\pm$ 0.006	0.575 $\pm$ 0.007	0.661$\pm$0.007	0.5	0.667$\pm$0.006	0.020
	0.30	0.680 $\pm$ 0.006	0.688 $\pm$ 0.006	0.742 $\pm$ 0.006	0.682 $\pm$ 0.006	0.680 $\pm$ 0.006	0.744$\pm$0.006	0.6	0.748$\pm$0.006	0.020
	0.40	0.712 $\pm$ 0.004	0.718 $\pm$ 0.003	0.800 $\pm$ 0.004	0.751 $\pm$ 0.004	0.746 $\pm$ 0.004	0.800$\pm$0.003	0.7	0.801$\pm$0.004	0.020

Table 2: As in Table 1, but here reporting the Balanced Accuracy metric. The proposed BL and PZ perform strongly throughout.

Dataset	η	CE	FL	GCE	ANL-CE	ANL-FL	BL	γ	PZ	c
MNIST	0.10	0.980 $\pm$ 0.002	0.971 $\pm$ 0.004	0.992 $\pm$ 0.001	0.992 $\pm$ 0.000	0.992 $\pm$ 0.000	0.993$\pm$0.001	0.4	0.992$\pm$0.001	0.040
	0.20	0.981 $\pm$ 0.001	0.974 $\pm$ 0.002	0.992 $\pm$ 0.000	0.991 $\pm$ 0.000	0.991 $\pm$ 0.000	0.993$\pm$0.000	0.5	0.993$\pm$0.000	0.020
	0.30	0.980 $\pm$ 0.001	0.971 $\pm$ 0.001	0.992 $\pm$ 0.000	0.991 $\pm$ 0.000	0.990 $\pm$ 0.000	0.993$\pm$0.000	0.5	0.993$\pm$0.000	0.020
	0.40	0.976 $\pm$ 0.001	0.966 $\pm$ 0.002	0.991 $\pm$ 0.000	0.990 $\pm$ 0.000	0.989 $\pm$ 0.000	0.992$\pm$0.000	0.5	0.991$\pm$0.000	0.040
Fashion-MNIST	0.10	0.942 $\pm$ 0.002	0.941 $\pm$ 0.002	0.949$\pm$0.002	0.943 $\pm$ 0.001	0.942 $\pm$ 0.001	0.951$\pm$0.001	0.5	0.948 $\pm$ 0.002	0.080
	0.20	0.943 $\pm$ 0.002	0.935 $\pm$ 0.002	0.949 $\pm$ 0.001	0.942 $\pm$ 0.001	0.940 $\pm$ 0.001	0.953$\pm$0.001	0.4	0.950$\pm$0.001	0.040
	0.30	0.939 $\pm$ 0.001	0.929 $\pm$ 0.001	0.948 $\pm$ 0.001	0.939 $\pm$ 0.001	0.938 $\pm$ 0.001	0.953$\pm$0.001	0.5	0.952$\pm$0.001	0.040
	0.40	0.932 $\pm$ 0.002	0.912 $\pm$ 0.027	0.946 $\pm$ 0.001	0.937 $\pm$ 0.001	0.936 $\pm$ 0.001	0.950$\pm$0.001	0.5	0.949$\pm$0.001	0.040
CIFAR-10	0.10	0.934 $\pm$ 0.002	0.931 $\pm$ 0.002	0.955 $\pm$ 0.002	0.946 $\pm$ 0.002	0.945 $\pm$ 0.002	0.956$\pm$0.001	0.7	0.956$\pm$0.002	0.060
	0.20	0.908 $\pm$ 0.002	0.912 $\pm$ 0.002	0.952 $\pm$ 0.003	0.941 $\pm$ 0.002	0.940 $\pm$ 0.002	0.952$\pm$0.002	0.7	0.953$\pm$0.002	0.080
	0.30	0.862 $\pm$ 0.003	0.879 $\pm$ 0.002	0.944 $\pm$ 0.002	0.934 $\pm$ 0.002	0.933 $\pm$ 0.002	0.945$\pm$0.002	0.7	0.945$\pm$0.002	0.080
	0.40	0.802 $\pm$ 0.003	0.834 $\pm$ 0.003	0.931 $\pm$ 0.002	0.925 $\pm$ 0.002	0.923 $\pm$ 0.002	0.933$\pm$0.002	0.8	0.934$\pm$0.002	0.100
CIFAR-100	0.10	0.833 $\pm$ 0.004	0.829 $\pm$ 0.004	0.853 $\pm$ 0.004	0.827 $\pm$ 0.005	0.825 $\pm$ 0.005	0.855$\pm$0.004	0.5	0.856$\pm$0.004	0.020
	0.20	0.815 $\pm$ 0.003	0.815 $\pm$ 0.003	0.850 $\pm$ 0.004	0.815 $\pm$ 0.005	0.812 $\pm$ 0.005	0.852$\pm$0.004	0.6	0.854$\pm$0.003	0.020
	0.30	0.785 $\pm$ 0.004	0.789 $\pm$ 0.005	0.844 $\pm$ 0.005	0.800 $\pm$ 0.005	0.797 $\pm$ 0.006	0.844$\pm$0.004	0.7	0.846$\pm$0.005	0.020
	0.40	0.746 $\pm$ 0.004	0.755 $\pm$ 0.003	0.833 $\pm$ 0.004	0.779 $\pm$ 0.005	0.774 $\pm$ 0.005	0.834$\pm$0.003	0.7	0.834$\pm$0.004	0.020

Table 1. The proposed loss functions are evaluated over a range of their parameters, γ and c ; however, only the best performing cases are shown for brevity. In nearly all cases, the use of the proposed BL and PZ losses result in better performance than with the baseline loss functions.

At low corruption rates ($\eta = 0.1$), FL occasionally performs well as consequence of it causing CL to detect more conservatively, favouring precision over recall compared to the proposed loss functions. At these low rates, the detection problem is highly imbalanced: high recall may be achieved while making few correct detections, but precision is dramatically reduced with relatively few false detections. In a context where detected examples are reviewed by experts, prioritizing recall over precision can be advantageous, as catching more errors generally outweighs the cost of incorrectly detecting a few more examples. Precision and recall are examined for BL and PZ, and compared to FL in Figure 3 of Section A.2, where it is found that by tuning the parameters of the proposed loss functions, recall is increased significantly, but at the cost of some precision. To capture a more reasonable objective at lower corruption rates, Balanced Accuracy (mean of sensitivity and specificity) is reported in Table 2, in which the proposed BL and PZ perform strongly throughout.

In the main experiment of Tables 1 and 2, the corrupted datasets include pre-existing, unaccounted for label errors. While the usage of several datasets demonstrated the broad applicability and efficacy of the proposed loss functions, an additional experiment on a dataset without pre-existing label errors would enable a more precise measure of label error detection performance. To achieve this, a version of CIFAR-100 with test samples having human-verified label errors removed [26] is used. The cleaned test set is artificially corrupted in the same manner as in the main experiment, and the Confident Learning error detection framework is once again applied; however, rather than using a k-fold approach over the entire dataset, model training is performed strictly on the non-cleaned training set, and evaluation is performed on the cleaned (and artificially corrupted) test set. Results

Table 3: Label error detections on the CIFAR-100 cleaned test set, averaged over three random trials, measured using the F1 score ($\pm\sigma$). Performance is improved relative to Table 1, especially at lower η and for PZ loss.

η	CE	FL	BL	PZ
0.1	0.576 $\pm$ 0.027	0.587$\pm$0.030	0.580 $\pm$ 0.024 ($\gamma = 0.3$)	0.590$\pm$0.020 ($c = 0.020$)
0.2	0.670 $\pm$ 0.013	0.671 $\pm$ 0.016	0.693$\pm$0.013 ($\gamma = 0.3$)	0.717$\pm$0.017 ($c = 0.015$)
0.3	0.697 $\pm$ 0.011	0.709 $\pm$ 0.006	0.775$\pm$0.011 ($\gamma = 0.5$)	0.780$\pm$0.013 ($c = 0.020$)
0.4	0.728 $\pm$ 0.005	0.729 $\pm$ 0.002	0.824$\pm$0.003 ($\gamma = 0.7$)	0.830$\pm$0.007 ($c = 0.025$)

Table 4: Label error detections on the CIFAR-10N and CIFAR-100N datasets with non-uniform corruption, averaged over five random trials, measured using the F1 score ($\pm\sigma$). The proposed BL and PZ loss functions continue to exceed baseline performance, with non-uniform corruption.

Dataset	CE	Focal	GCE	ANL-CE	ANL-FL	BL	γ	PZ	c
CIFAR-10N, $\eta \approx 0.1$	0.662 $\pm$ 0.004	0.634 $\pm$ 0.002	0.698 $\pm$ 0.002	0.652 $\pm$ 0.006	0.647 $\pm$ 0.003	0.698 $\pm$ 0.004	0.7	0.704 $\pm$ 0.006	0.08
CIFAR-100N, $\eta \approx 0.4$	0.700 $\pm$ 0.002	0.684 $\pm$ 0.002	0.741 $\pm$ 0.002	0.721 $\pm$ 0.002	0.721 $\pm$ 0.003	0.742 $\pm$ 0.003	0.8	0.745 $\pm$ 0.003	0.10
CIFAR-10N, $\eta \approx 0.4$	0.768 $\pm$ 0.001	0.769 $\pm$ 0.002	0.837 $\pm$ 0.002	0.865 $\pm$ 0.001	0.865 $\pm$ 0.002	0.867 $\pm$ 0.001	0.9	0.866 $\pm$ 0.001	0.15

averaged over three seeds are shown in Table 3 and form a direct comparison to the CIFAR-100 experiments shown in Table 1. The performance of all loss functions is improved, especially at lower corruption rates (improvement of nearly 0.1 in F1 score at $\eta = 0.1$). This improvement at $\eta = 0.1$ matches expectations given that the proportion of pre-existing label errors is more significant at lower artificial corruption rates. The most dramatic improvement is seen for PZ loss for all η , which performs best in all cases.

6.2 Ablation Studies

Non-uniform Corruption The main experiment included only *uniform* artificial corruption. To better understand the performance of the proposed loss functions under the presence of *non-uniform* label errors, a study involving the CIFAR-10N and CIFAR-100N [38] datasets with real-world human annotation errors is performed. This study follows the same format as the main experiment. Results averaged over five random seeds are shown in Table 4 against baseline loss functions, and for three variations of the CIFAR-N datasets, as described in Section 5. Rows are ordered with increasing corruption rate. Notice that the proposed loss functions are best (and second best) performing in all cases, and the optimal parameter setting tends to increase with corruption rate, congruent with the results of ablation Figure 4 in Section A.2.

Alternative Label Error Detection Framework The main experiment used Confident Learning as the label error detection framework. To demonstrate the effectiveness and versatility of the proposed loss functions when used with other frameworks, an ablation study using the AUM [31] framework is performed on the CIFAR-100 dataset. Resulting label error detection F1 scores averaged over three random seeds are summarized in Table 5 and are directly comparable to the CIFAR-100 results of Table 1. The label error detection performance exceeds that of Confident Learning, and the proposed loss functions improve upon the baseline.

Training Dynamics To verify that the proposed loss functions affect gradients during training as they are designed to do, the gradient of loss, $\partial\mathcal{L}/\partial p_{\tilde{y}}$, is monitored during training for corrupted and

Table 5: Label error detections using the AUM framework and with the CIFAR-100 dataset, averaged over three random trials, measured using the F1 score ($\pm\sigma$). The proposed losses continue to perform well using AUM, in addition to with CL (Tables 1 and 2).

η	Loss Function		
	CE	BL ($\gamma = 0.05$)	PZ ($c = 0.025$)
0.1	0.7452 $\pm$ 0.0071	0.7591 $\pm$ 0.0180	0.8817 $\pm$ 0.0074
0.2	0.8638 $\pm$ 0.0042	0.8692 $\pm$ 0.0032	0.8903 $\pm$ 0.0085
0.3	0.9019 $\pm$ 0.0021	0.9076 $\pm$ 0.0031	0.8914 $\pm$ 0.0044
0.4	0.9197 $\pm$ 0.0024	0.9261 $\pm$ 0.0005	0.8856 $\pm$ 0.0125

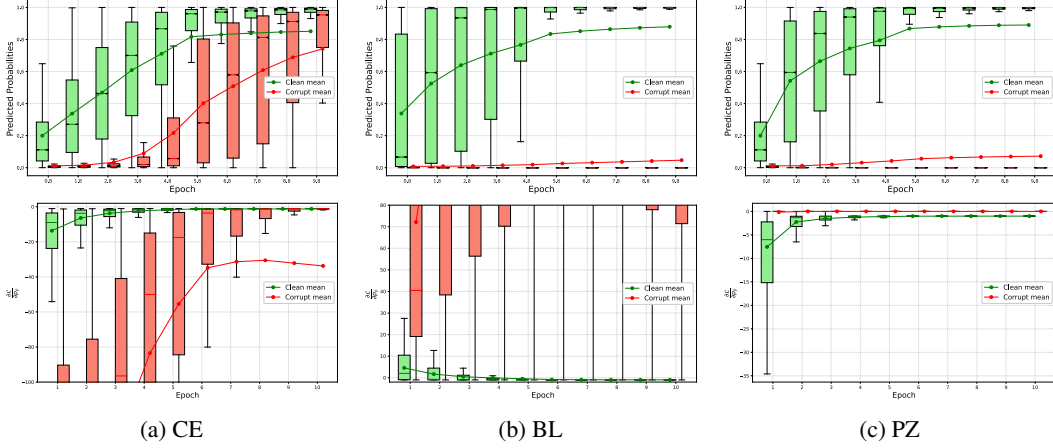


Figure 2: Predicted probability (top) and Gradient (bottom) distributions per epoch, for CIFAR-100 at $\eta = 0.4$. With all loss functions, the predicted probabilities of corrupt data (red) are less than those of clean data (green). With CE, the gradients of corrupt data are large and negative, providing a strong signal to the optimizer to fit to these corrupt data, whereas with BL, gradients of corrupt data are large and *positive*, steering away from these corrupt data, and with PZ, gradients of corrupt data are nearly all at zero, imparting no impact on training. A version of this figure without truncated vertical limits can be found in Figure 7 of Section A.3.

clean samples. Box and whisker plots in Figure 2 indicate the distribution of predicted probabilities (top) and gradients (bottom) for corrupt (red) and clean (green) data at each epoch during training. CIFAR-100 with $\eta = 0.4$ is used for this study. With all loss functions, the predicted probabilities of corrupt data are less than those of clean data; however, the difference is far greater with the proposed losses. With CE, the gradients of corrupt data are large and negative, providing a strong signal to the optimizer to fit to these corrupt data, whereas with BL, gradients of corrupt data are mainly large but *positive*, steering the training away from these corrupt data, and with PZ, gradients of corrupt data are nearly all at zero, imparting no impact on training. Notice that with BL, the gradients for clean data during the first epochs are mostly *positive*, but quickly settle to a moderate *negative* value, associated with predicted probabilities close to 1.0 (as can be seen in the top row).

Further Studies Additional investigations and ablation studies are provided in Section A.2, including a study observing the precision-recall trade-off and their impact on F1 score, in Figure 3, a study investigating the impacts of the proposed loss function parameters, in Figure 4, an ablation study on the effects of the delay parameter, d , in Figure 5, and an additional study observing differences in distributions of the AUM statistic for corrupt and clean samples, in Figure 6.

7 Conclusion and Limitations

This work provides an in-depth empirical evaluation of two recently proposed loss functions, Blurry Loss and Piecewise-zero Loss, for improving model robustness and the downstream detection of label errors. By de-weighting or ignoring samples likely to be mislabelled, these losses enhance the performance of label error detection frameworks across a variety of artificial corruption settings. Extensive experiments and ablation studies demonstrate their broad applicability and competitive performance relative to existing robust loss functions. One limitation of the experiments is that parameter sweeps were not performed for the existing robust loss functions; however, their recommended settings for each dataset were used wherever possible. Similarly, there is no universally optimal parameter setting for the proposed loss functions. These findings support the use of such loss functions as practical tools for improving data quality and model reliability in supervised learning scenarios fraught with label errors.

Acknowledgments

This research was enabled in part by support provided by Calcul Québec (calculquebec.ca) and the Digital Research Alliance of Canada (alliancecan.ca).

We acknowledge the support of the Government of Canada’s New Frontiers in Research Fund (NFRF), [NFRFT-2020-00073], and the support of the Natural Sciences and Engineering Research Council of Canada (NSERC) via NSERC-CGS D.

Nous remercions le Fonds Nouvelles Frontières en Recherche du gouvernement du Canada de son soutien (FNFR), [FNFR-2020-00073], et le soutien du Conseil de Recherches en Sciences Naturelles et en Génie du Canada (CRSNG), CRSNG-BESC D.



Natural Sciences and Engineering
Research Council of Canada

Conseil de recherches en sciences
naturelles et en génie du Canada

Canada

References

- [1] George EP Box and David R Cox. An analysis of transformations. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 26(2):211–243, 1964.
- [2] Kay Henning Brodersen, Cheng Soon Ong, Klaas Enno Stephan, and Joachim M. Buhmann. The balanced accuracy and its posterior distribution. In *2010 20th International Conference on Pattern Recognition*, pages 3121–3124, 2010. doi: 10.1109/ICPR.2010.764.
- [3] Jacob Cohen. *Statistical power analysis for the behavioral sciences*. routledge, 2013.
- [4] Li Deng. The mnist database of handwritten digit images for machine learning research. *IEEE signal processing magazine*, 29(6):141–142, 2012.
- [5] Camille Garcin, Alexis Joly, Pierre Bonnet, Jean-Christophe Lombardo, Antoine Affouard, Mathias Chouet, Maximilien Servajean, Titouan Lorieul, and Joseph Salmon. Pl@ntNet-300K: a plant image dataset with high label ambiguity and a long-tailed distribution. In *Conference on Neural Information Processing Systems (NeurIPS) Datasets and Benchmarks*, 2021.
- [6] Zahra Gharaee, ZeMing Gong, Nicholas Pellegrino, Iuliia Zarubiieva, Joakim Bruslund Haurum, Scott Lowe, Jaclyn McKeown, Chris Ho, Joschka McLeod, Yi-Yun Wei, et al. A step towards worldwide biodiversity assessment: The bioscan-1m insect dataset. *Advances in Neural Information Processing Systems*, 36, 2024.
- [7] Zahra Gharaee, Scott C Lowe, ZeMing Gong, Pablo Millan Arias, Nicholas Pellegrino, Austin T Wang, Joakim Bruslund Haurum, Iuliia Eyriay, Lila Kari, Dirk Steinke, et al. Bioscan-5m: A multimodal dataset for insect biodiversity. *Advances in Neural Information Processing Systems*, 37:36285–36313, 2024.
- [8] Aritra Ghosh, Naresh Manwani, and PS Sastry. Making risk minimization tolerant to label noise. *Neuro-computing*, 160:93–107, 2015.
- [9] Aritra Ghosh, Himanshu Kumar, and P Shanti Sastry. Robust loss functions under label noise for deep neural networks. In *Proceedings of the AAAI conference on artificial intelligence*, volume 31, 2017.
- [10] Irving John Good. Rational decisions. *Journal of the Royal Statistical Society: Series B (Methodological)*, 14(1):107–114, 1952.
- [11] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 770–778, 2016.
- [12] Wei He, Kai Han, Ying Nie, Chengcheng Wang, and Yunhe Wang. Species196: A one-million semi-supervised dataset for fine-grained species recognition. *Advances in Neural Information Processing Systems*, 36, 2024.
- [13] Lu Jiang, Zhengyuan Zhou, Thomas Leung, Li-Jia Li, and Li Fei-Fei. MentorNet: Learning data-driven curriculum for very deep neural networks on corrupted labels. In *International conference on machine learning*, pages 2304–2313. PMLR, 2018.
- [14] Shivam S Kadam, Amol C Adamuthe, and Ashwini B Patil. Cnn model for image classification on mnist and fashion-mnist dataset. *Journal of scientific research*, 64(2):374–384, 2020.

- [15] Leonid V Kantorovich. Mathematical methods of organizing and planning production. *Management science*, 6(4):366–422, 1960.
- [16] Youngdong Kim, Junho Yim, Juseung Yun, and Junmo Kim. Nlnl: Negative learning for noisy labels. In *Proceedings of the IEEE/CVF international conference on computer vision*, pages 101–110, 2019.
- [17] Diederik P Kingma. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*, 2014.
- [18] Alex Krizhevsky and Geoffrey Hinton. Learning multiple layers of features from tiny images. Technical Report 0, University of Toronto, Toronto, Ontario, 2009. URL <https://www.cs.toronto.edu/~kriz/learning-features-2009-TR.pdf>.
- [19] T Lin. Focal loss for dense object detection. *arXiv preprint arXiv:1708.02002*, 2017.
- [20] Yueming Lyu and Ivor W Tsang. Curriculum loss: Robust learning and generalization against label corruption. In *International Conference on Learning Representations*, 2019.
- [21] Xingjun Ma, Hanxun Huang, Yisen Wang, Simone Romano, Sarah Erfani, and James Bailey. Normalized loss functions for deep learning with noisy labels. In *International conference on machine learning*, pages 6543–6553. PMLR, 2020.
- [22] Aditya Krishna Menon, Ankit Singh Rawat, Sashank J Reddi, and Sanjiv Kumar. Can gradient clipping mitigate label noise? In *International conference on learning representations*, 2020.
- [23] Nagarajan Natarajan, Inderjit S Dhillon, Pradeep K Ravikumar, and Ambuj Tewari. Learning with noisy labels. *Advances in neural information processing systems*, 26, 2013.
- [24] H. Nguyen, T. Truong, X. Nguyen, A. Dowling, X. Li, and K. Luu. Insect-Foundation: A foundation model and large-scale 1M dataset for visual insect understanding. In *IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 21945–21955, Los Alamitos, CA, USA, Jun 2024. IEEE Computer Society. doi: 10.1109/CVPR52733.2024.02072.
- [25] Curtis Northcutt, Lu Jiang, and Isaac Chuang. Confident learning: Estimating uncertainty in dataset labels. *Journal of Artificial Intelligence Research*, 70:1373–1411, 2021.
- [26] Curtis G Northcutt, Anish Athalye, and Jonas Mueller. Pervasive label errors in test sets destabilize machine learning benchmarks. *Advances in Neural Information Processing Systems*, 2021.
- [27] Tianyu Pang, Chao Du, Yinpeng Dong, and Jun Zhu. Towards robust detection of adversarial examples. *Advances in neural information processing systems*, 31, 2018.
- [28] Giorgio Patrini, Alessandro Rozza, Aditya Krishna Menon, Richard Nock, and Lizhen Qu. Making deep neural networks robust to label noise: A loss correction approach. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 1944–1952, 2017.
- [29] Nicholas Pellegrino, Nolen Zhao, and Paul Fieguth. The effects of label errors in training data on model performance and overfitting. *Journal of Computational Vision and Imaging Systems*, 9(1):26–29, 2023.
- [30] Nicholas Pellegrino, David Szczecina, and Paul Fieguth. Loss functions robust to the presence of label errors. *Journal of Computational Vision and Imaging Systems*, 10(1):24–29, 2024.
- [31] Geoff Pleiss, Tianyi Zhang, Ethan Elenberg, and Kilian Q Weinberger. Identifying mislabeled data using the area under the margin ranking. *Advances in Neural Information Processing Systems*, 33:17044–17056, 2020.
- [32] Hwanjun Song, Minseok Kim, and Jae-Gil Lee. Selfie: Refurbishing unclean samples for robust deep learning. In *International conference on machine learning*, pages 5907–5915. PMLR, 2019.
- [33] Hwanjun Song, Minseok Kim, Dongmin Park, Yooju Shin, and Jae-Gil Lee. Learning from noisy labels with deep neural networks: A survey. *IEEE transactions on neural networks and learning systems*, 34(11): 8135–8153, 2022.
- [34] Grant Van Horn, Elijah Cole, Sara Beery, Kimberly Wilber, Serge Belongie, and Oisín MacAodha. Benchmarking representation learning for natural world image collections. In *IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 12879–12888, 2021. doi: 10.1109/CVPR46437.2021.01269.
- [35] Cornelius Joost Van Rijsbergen. Information retrieval. 2nd. newton, ma, 1979.

- [36] Leonid Nisonovich Vaserstein. Markov processes over denumerable products of spaces, describing large systems of automata. *Problemy Peredachi Informatsii*, 5(3):64–72, 1969.
- [37] Yisen Wang, Xingjun Ma, Zaiyi Chen, Yuan Luo, Jinfeng Yi, and James Bailey. Symmetric cross entropy for robust learning with noisy labels. In *Proceedings of the IEEE/CVF international conference on computer vision*, pages 322–330, 2019.
- [38] Jiaheng Wei, Zhaowei Zhu, Hao Cheng, Tongliang Liu, Gang Niu, and Yang Liu. Learning with noisy labels revisited: A study using real-world human annotations. In *International Conference on Learning Representations*, 2022. URL <https://openreview.net/forum?id=TBWA6PLJZQm>.
- [39] Xiaoping Wu, Chi Zhan, Yu-Kun Lai, Ming-Ming Cheng, and Jufeng Yang. IP102: A large-scale benchmark dataset for insect pest recognition. In *IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 8779–8788, 2019. doi: 10.1109/CVPR.2019.00899.
- [40] Han Xiao, Kashif Rasul, and Roland Vollgraf. Fashion-mnist: a novel image dataset for benchmarking machine learning algorithms. *arXiv preprint arXiv:1708.07747*, 2017.
- [41] Xichen Ye, Xiaoqiang Li, Tong Liu, Yan Sun, Weiqin Tong, et al. Active negative loss functions for learning with noisy labels. *Advances in Neural Information Processing Systems*, 36:6917–6940, 2023.
- [42] Chiyuan Zhang, Samy Bengio, Moritz Hardt, Benjamin Recht, and Oriol Vinyals. Understanding deep learning requires rethinking generalization. *arXiv preprint arXiv:1611.03530*, 2016.
- [43] Zhilu Zhang and Mert Sabuncu. Generalized cross entropy loss for training deep neural networks with noisy labels. *Advances in neural information processing systems*, 31, 2018.
- [44] Xiong Zhou, Xianming Liu, Junjun Jiang, Xin Gao, and Xiangyang Ji. Asymmetric loss functions for learning with noisy labels. In *International conference on machine learning*, pages 12846–12856. PMLR, 2021.

A Technical Appendices and Supplementary Material

A.1 Additional Experimental Details

Further information building upon the experimental details of Section 5 is given here.

Models and Optimizers For MNIST and Fashion-MNIST, we adopt a minimal convolutional neural network consisting of two convolutional layers followed by two fully connected layers interleaved with dropout, totalling approximately 1.2 million trainable parameters. Under a ten-epoch training schedule, this architecture reliably exceeds 99% test accuracy (near state-of-the-art) [14].

For CIFAR-10, CIFAR-100, and their noisy variants (CIFAR-10N/CIFAR-100N), we fine-tune an ImageNet-pretrained ResNet-34 [11] modified only by replacing its initial 7×7 , stride-2 convolution with a 3×3 , stride-1 convolution and removing the following max-pool layer, preserving spatial resolution on 32×32 inputs. We train for ten epochs under identical conditions using the Adam optimizer [17] with an initial learning rate of 1×10^{-4} ; a scheduler reduces the learning rate by a factor of 0.1 at epoch 5 for CIFAR-10 (factor 0.2 for CIFAR-100). Under this setup, the model attains 93.5% test accuracy on CIFAR-10 and 74.5% on CIFAR-100.

All experiments utilize the Adam optimizer, leveraging its adaptive learning capabilities without requiring additional tuning.

Compute Resource Requirements All experiments were conducted on a compute cluster using an NVIDIA Tesla V100-SXM2-16GB GPU. The main Confident Learning experiments required 6 GB of memory, and each loss function and parameter combination required on average 10 minutes (5 minutes for MNIST and Fashion-MNIST, 15 minutes for CIFAR10 and CIFAR100). Experiments on the cleaned CIFAR-100 test set followed the same process, requiring only 3 GB memory. When additionally storing the per-epoch gradient information, for MNIST 3 GB of memory and 1 hour was required, and 6 GB of memory and 4 hours for CIFAR100. The CIFAR-N experiments required 15 minutes and 6 GB of memory. The AUM experiments required 7 GB of memory as well as 2.5 hours for each loss function and parameter test conducted. Our complete set of experiments required approximately 1800 hours of compute, with an additional ~ 400 hours for preliminary experimentation.

The Confident Learning experiments were conducted using the Python library `cleanlab` v2.5.0, from <https://github.com/cleanlab/cleanlab>, as introduced in Northcutt et al. [25] under an AGPL-3.0 license. The AUM experiments were conducted using the python library `aum` v1.0.2, from <https://github.com/asappresearch/aum> as introduced in Pleiss et al. [31] under an MIT license.

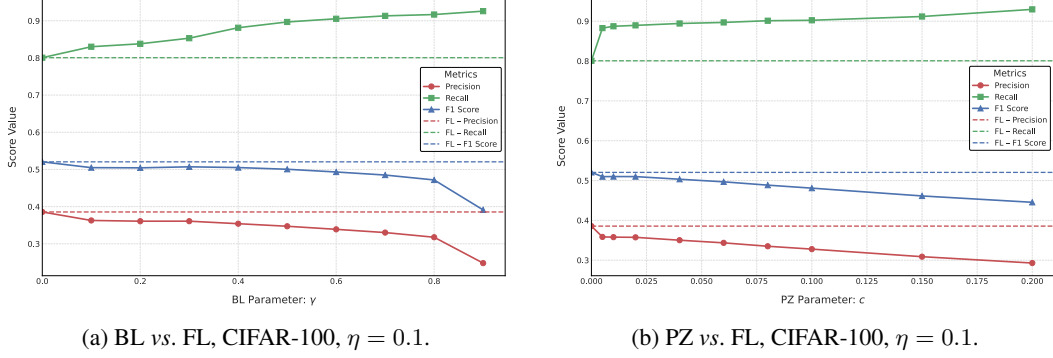


Figure 3: FL (dotted lines) results in better F1 scores than BL (a) or PZ (b) (solid lines) at $\eta = 0.1$ on the CIFAR-100 dataset. This is a result of conservative detection, whereby fewer samples are detected and precision is maintained. Through tuning the parameters of the proposed loss functions, performance can be made fairly similar to that of FL, or recall can be increased dramatically while only marginally reducing precision, which may be considered worthwhile at low corruption rates.

A.2 Additional Experiments and Studies

Precision-Recall Trade-off in F1 Score Precision and recall (the two harmonic components of the F1 score) are explored for BL and PZ, and compared to FL in Figure 3, for the CIFAR-100 dataset at $\eta = 0.1$, a case selected since the F1 score of FL exceeds that of BL and PZ. Observe that by tuning the parameters of the proposed loss functions, recall is increased significantly, but at the cost of some precision. Because F1 is a harmonic mean, the small decrease in precision dominates the larger increase in recall, and the score is decreased. Again, this does not necessarily indicate less desirable behaviour at lower corruption rates, since having improved recall may be worthwhile at the expense of having relatively few false-positives.

Loss Function Parameters To better understand the parameters of the proposed loss functions, an ablation study on these parameters is performed using a similar experimental setup as the main experiment. For each dataset and artificial corruption rate, the optimal parameter setting is identified. For these experiments, the delay, d , is set to 0 for BL and 1 for PZ loss. Results for MNIST and CIFAR-100 datasets are shown in Figure 4 as a heatmap of F1 score, along with a comparison to the baseline loss functions. For each corruption rate (row), the best performing configuration is indicated with bolded yellow text. To illustrate the trend in optimal parameter settings for the proposed loss functions, yellow boxes mark the best performing settings. Observe that optimal settings of both γ and cutoff, c , increase with increased corruption rate. This trend is consistent across datasets; however, more complicated / difficult datasets, such as CIFAR-100 relative to MNIST, seem also to demand greater settings of γ (but not c). Note that larger settings for γ and c correspond to the proposed loss functions ‘ignoring’ a greater range of predicted probabilities.

Delay Delay, d , is introduced to allow the model to generalize prior to training with the proposed loss functions. This is especially important for PZ loss, which, at the initial stages of training, outright ‘ignores’ a large proportion of the dataset. To better understand how best to set the delay, an ablation study on this parameter is performed. Again, label errors are detected and performance is measured with F1 score, while d is varied. This study was performed across all the main experiment datasets and corruption rates, with fairly consistent results.

Results for the most complex dataset, CIFAR-100, are given in Figure 5, for corruption rates ranging from $\eta = 0.1$ to 0.4. Trends for various cutoff, c , settings are plotted. Observe that there is a large jump in F1 from $d = 0$ to $d = 1$ (for all η), indicating that some delay is highly beneficial and allowing the model to somewhat generalize early in training is necessary for good application of PZ loss. Additionally, the best performance is achieved for $d = 1$ (for an appropriate setting of c), and F1 scores tend to towards the CE baseline (poorer performance) as delay is further increased. This indicates that having a greater number of epochs trained with PZ loss improves label error detection once the model is sufficiently generalized (one epoch with CE), and perhaps even that detrimental fitting to the erroneous data occurs with further training with CE. Furthermore, notice that the degree of downwards trending relates to the corruption rate, indicating that high corruption rates necessitate a greater need for training with PZ loss to achieve high label error detection performance.

Although not shown here, a similar preliminary study was performed for BL, which indicated that having no delay was always best, and thus delay is suggested only for PZ loss. This is likely a result of PZ loss being far more harsh than BL (outright ignoring of data rather than simply having positive gradients for samples with low predicted probabilities).

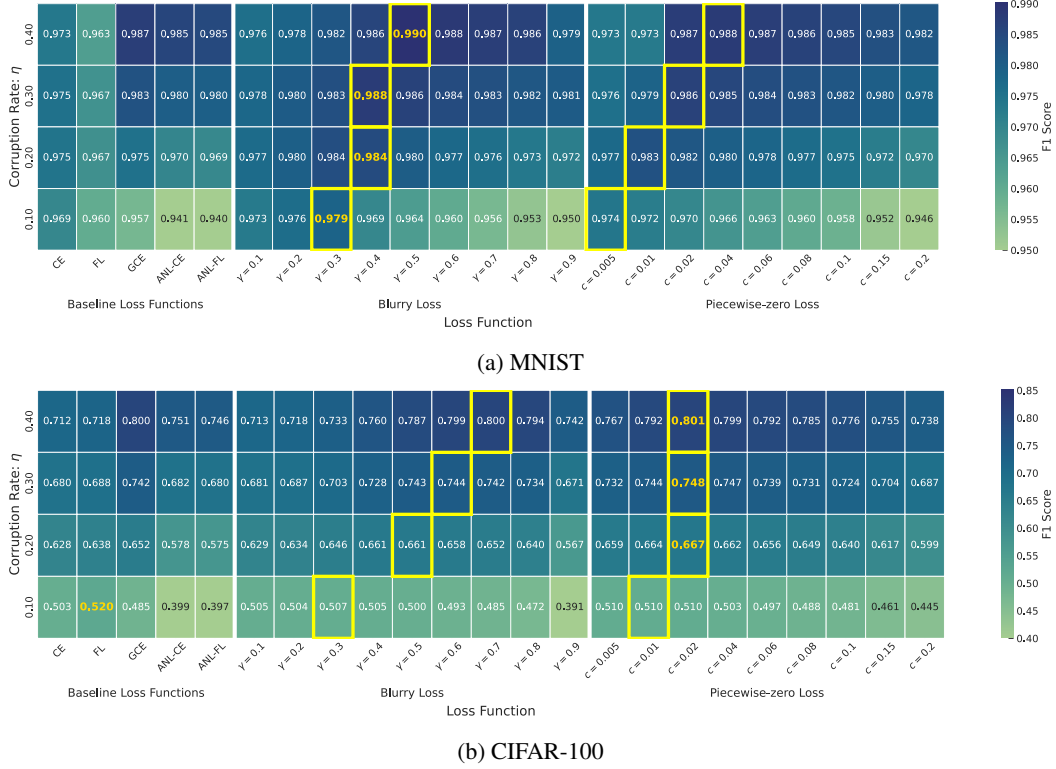


Figure 4: Heatmap showing variation in F1 score as loss function parameters are varied, for MNIST (a) and CIFAR-100 (b). In the left blocks, separated by a white line, results for baseline loss functions are shown. Two separate blocks are shown for the proposed loss functions, plotted over a range in their parameters. Best results are indicated with bolded yellow text. To illustrate the trend in optimal parameter settings for the proposed loss functions, yellow boxes mark the best performing parameter setting for each row. Note that delay, d , is set to 0 for BL and 1 for PZ loss.

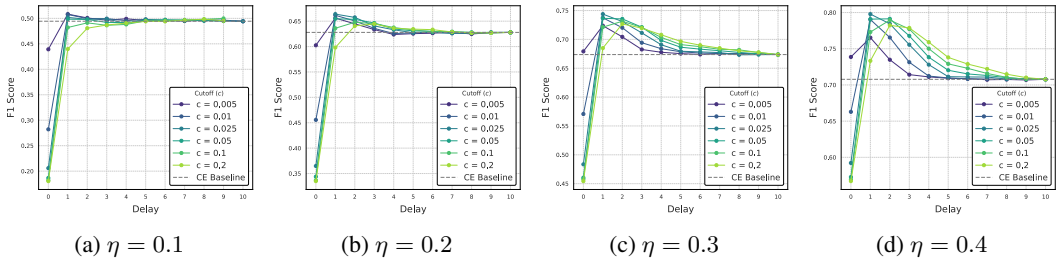


Figure 5: Detection F1 score vs. Delay, d , for PZ loss on CIFAR-100 at η ranging from 0.1 to 0.4 in panels (a) to (d). Observe that the best performance comes at $d = 1$ in all cases, for an appropriate setting of the cutoff, c .

Training Dynamics As was explained in Section 3, the AUM statistic [31] measures the difference between the logit values for a training example’s as-labelled class and its highest other class, averaged over the course of training. Correctly labelled samples tend to exhibit positive AUM values, while mislabelled samples tend to have significantly lower AUM values due to conflicting gradient updates for samples of the same class. To better understand how the proposed loss functions affect the training dynamics of corrupt vs. clean data, the AUM is measured for all samples, and the AUM distributions for corrupt and clean data are compared. Figure 6 illustrates the AUM distributions for training with CE, BL, and PZ loss on the CIFAR-100 dataset with artificial corruption at $\eta = 0.1$ (following the same method used throughout this paper). The mean of each distribution is indicated using a dashed vertical line. To quantitatively understand the impacts of training using the proposed loss functions, two measures of dissimilarity in distributions are computed: Cohen’s d [3] (Effect size) and Wasserstein distance [15, 36], both of which are sensitive to differences in distribution position and width. Note

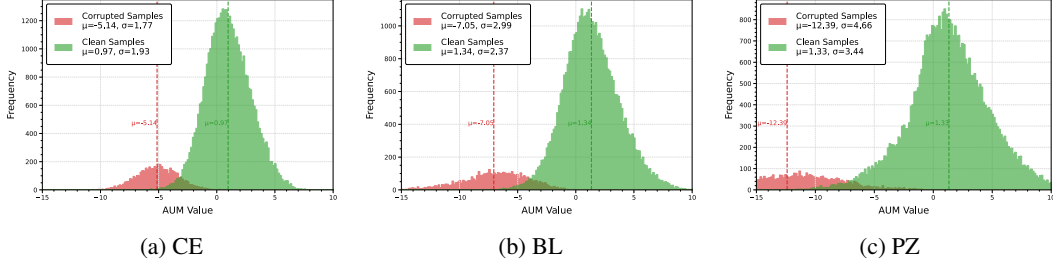


Figure 6: AUM distributions for artificially corrupted and clean samples in CIFAR-100 with corruption at $\eta = 0.1$ for training with CE (a), BL (b), and PZ loss (c). Greater separation in distributions indicates a greater robustness of the loss function to mislabelled samples during training. Separation is measured using Cohen’s d {3.20 (CE) < 3.43 (BL) < 3.82 (PZ)} and Wasserstein distance {6.11 (CE) < 8.40 (BL) < 13.72 (PZ)}.

that an ideal robust training scheme would result in a large difference in the AUM distributions of corrupt and clean samples.

Under CE training (Figure 6a), clean and corrupted samples exhibit distinct but moderately overlapping distributions. Although separation exists, the overlap reflects CE’s tendency to fit both clean and noisy labels, leading to partial memorization of mislabelled samples.

Training with Blurry Loss (Figure 6b), produced a markedly improved separation. Quantitatively, a Cohen’s d of 3.43 (vs. 3.20 for CE), and Wasserstein distance of 8.40 (vs. 6.11 for CE) were measured, indicating improved separation of corrupt and clean data. This increased separation of the distributions supports our intention and hypothesis that BL de-weights hard-to-classify (likely mislabelled) samples, reducing the extent to which the model fits to label errors. Consequently, the discrimination between clean and noisy samples becomes more pronounced, enhancing the effectiveness of AUM-based label error detection.

The most dramatic effect was observed when using Piecewise-zero Loss (Figure 6c). Here, Cohen’s d increased to 3.82 (vs. 3.20 for CE), and the Wasserstein distance to 13.72 (vs. 6.11 for CE), indicating greater separation than both CE and BL. The sharp cutoff mechanism of PZ effectively prevents low-confidence samples from contributing to weight updates during training, almost eliminating their influence. As a result, mislabelled examples tend consistently to have highly negative AUM values, leading to the clearest distributional separation among all loss functions tested. Indeed, the F1 score notably increased by more than 13% relative to CE, as was shown in Table 5.

Both BL and PZ loss functions substantially improve the margin-based separability of clean versus mislabelled samples compared to CE. These results validate our design objectives: the proposed loss functions mitigate the harmful fitting to noisy labels by selective de-weighting or exclusion. This enhanced separability directly translates into improved performance for label error detection frameworks like AUM and Confident Learning.

A.3 Expanded Versions of Results Shown in Main Paper

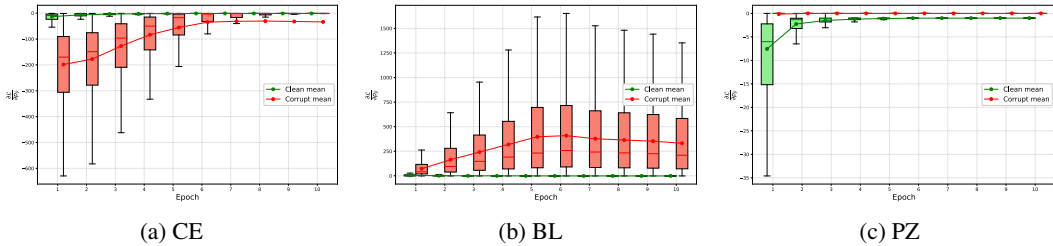


Figure 7: A version of Figure 2 without truncated vertical limits. Distributions of gradients of corrupt (red) and clean (green) data are plotted for each epoch for CIFAR-100 at $\eta = 0.4$.