

Setting

Supervised classification with instance space $\mathcal{X}$ and label space $\mathcal{Y}$.

Training data

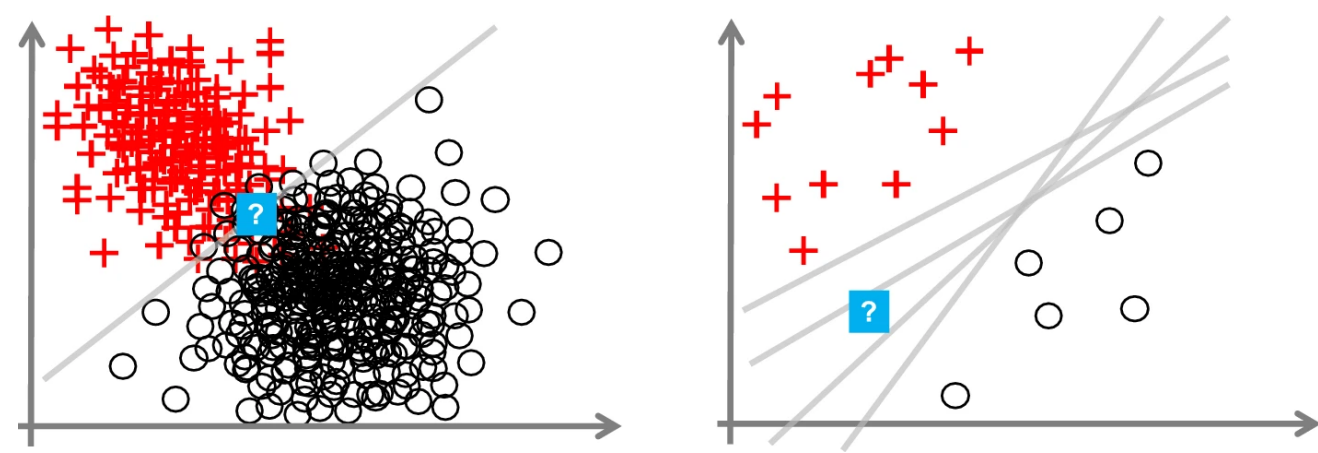
$$\mathcal{D} = \{(\mathbf{x}_i, y_i)\}_{i=1}^N \in (\mathcal{X} \times \mathcal{Y})^N.$$

Hypothesis space

$$\mathcal{H} = \{h : \mathcal{X} \rightarrow \mathbb{P}(\mathcal{Y})\}.$$

Given an instance $\mathbf{x}$, we denote $h(\mathbf{x}) = \theta$.

Types of Uncertainty [1]



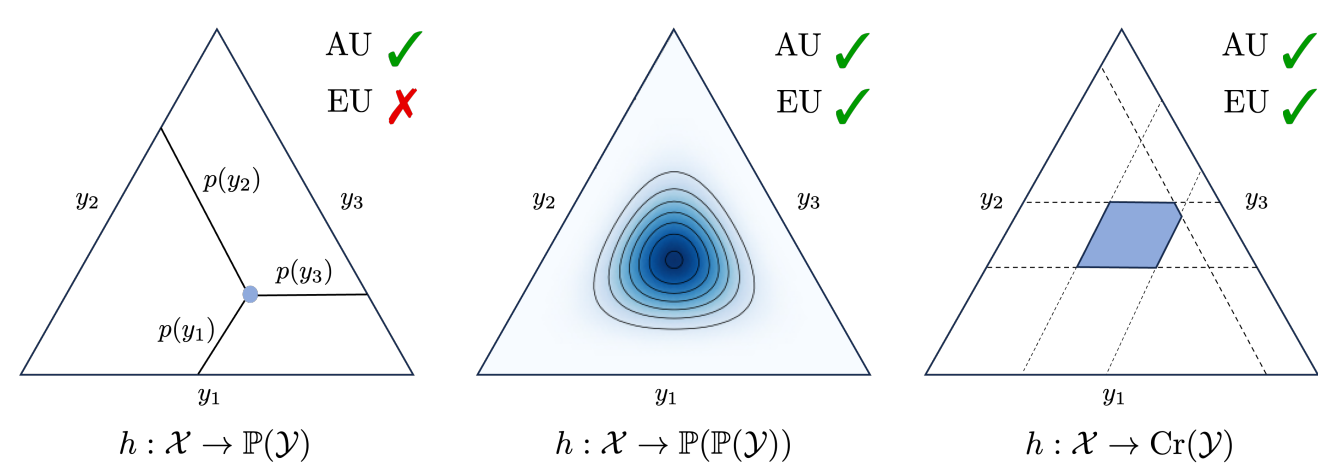
Aleatoric Uncertainty

Refers to the variability of the outcome due to the inherent **randomness** of the data and is therefore **irreducible**.

Epistemic Uncertainty

Refers to uncertainty caused by a **lack of knowledge** (epistemic state of the learner) and is **reducible** by adding information.

Uncertainty Representation



A credal set C is a (closed and convex) set of probability distributions on $\mathcal{Y}$.

$$TU(C) = AU(C) + EU(C).$$

Various decompositions exist, i.a. based on

$$GH(C) := \sum_{A \subseteq \mathcal{Y}} m_Q(A) \log(|A|)$$

gen. Hartley [2] and Shannon entropy S .

$$\underbrace{S^*(C)}_{TU} = \underbrace{(S^*(C) - GH(C))}_{AU} + \underbrace{GH(C)}_{EU},$$

$$\underbrace{S^*(C)}_{TU} = \underbrace{S_*(C)}_{AU} + \underbrace{(S^*(C) - S_*(C))}_{EU},$$

where

$$S^*(C) := \max_{\theta \in C} S(\theta), \quad S_*(C) := \min_{\theta \in C} S(\theta).$$

A credal approach

Learning Credal Predictors

Relative likelihood given training data $\mathcal{D}$

$$\mathcal{L}_{\mathcal{H}}(h) := \frac{L(h)}{L(\hat{h})},$$

where $L(h) = \prod_{i=1}^N p(y_i | h, \mathbf{x})$ and $\hat{h}$ the (empirical) maximum likelihood predictor.

Given $\mathbf{x} \in \mathcal{X}$ and $Q \subseteq \mathcal{H}$ (plausible hypotheses, i.e. ensemble members):

$$C_{\alpha} := \{h(\mathbf{x}) \in Q \mid \mathcal{L}_{\mathcal{H}}(h) \geq \alpha\},$$

with $\alpha \in [0, 1]$.

Uncertainty Quantification

Epistemic uncertainty

$$EU(C) := \max_{\theta, \theta' \in C} D_{\ell}(\theta, \theta'),$$

with the ℓ -divergence

$$D_{\ell}(\theta, \theta') := \mathbb{E}_{Y \sim \theta} \{\ell(\theta', Y) - \ell(\theta, Y)\},$$

the excess loss of predicting θ' , while the ground truth is θ .

Aleatoric uncertainty

$$\underline{AU}(C) := \min_{\theta \in C} H_{\ell}(\theta),$$

$$\overline{AU}(C) := \max_{\theta \in C} H_{\ell}(\theta),$$

with H_{ℓ} the ℓ -entropy of θ given by

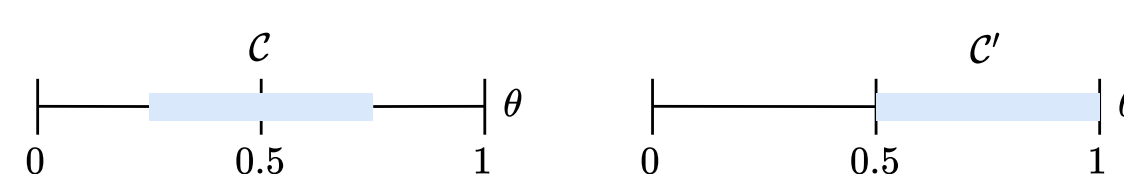
$$H_{\ell}(\theta) := \mathbb{E}_{Y \sim \theta} \ell(\theta, Y),$$

the irreducible loss of ground truth θ .

Here, ℓ can be any loss. We consider (strictly) proper scoring rules [3].

Loss	Aleatoric (upper\lower)	Epistemic
<i>log</i>	$\sup_{\theta \in C} \inf_{\theta' \in C} S(\theta)$	$\max_{\theta, \theta' \in C} D_{KL}(\theta' \parallel \theta)$
<i>Brier</i>	$\sup_{\theta \in C} \inf_{\theta' \in C} 1 - \sum_{k=1}^K \theta_k^2$	$\max_{\theta, \theta' \in C} \sum_{k=1}^K (\theta_k - \theta'_k)^2$
<i>spherical</i>	$\sup_{\theta \in C} \inf_{\theta' \in C} 1 - \ \theta\ _2$	$\max_{\theta, \theta' \in C} \ \theta'\ _2 - \sum_{k=1}^K \theta_k \theta'_k / \ \theta'\ _2$
<i>zero-one</i>	$\sup_{\theta \in C} \inf_{\theta' \in C} 1 - \max \theta_k$	$\max_{\theta, \theta' \in C} \max_k \theta'_k - \theta'_{k=\arg \max \theta_k}$

Different losses account for different uncertainties [4].

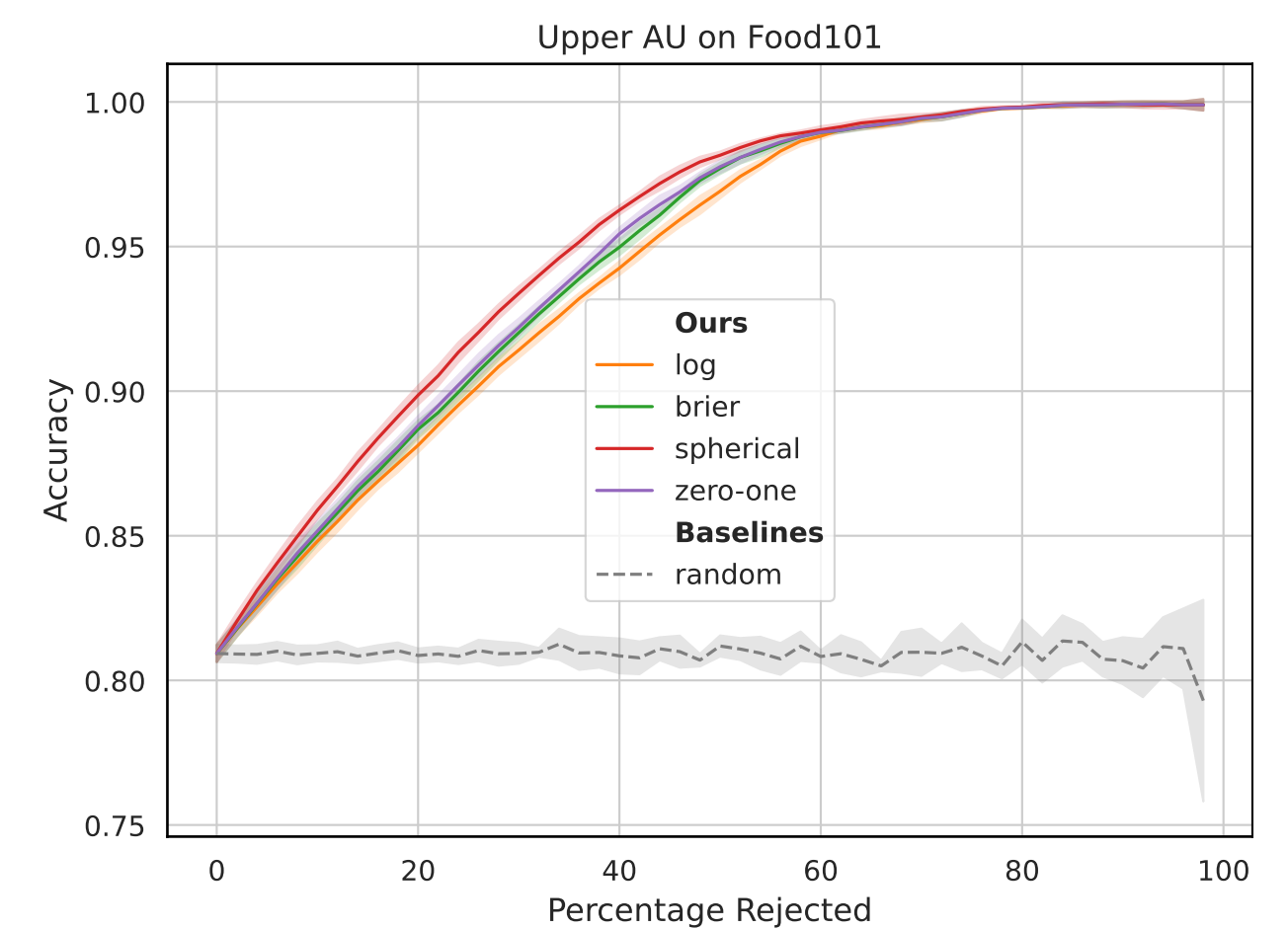


Left: decision uncertainty. Right: no decision uncertainty.

Results

Accuracy-Rejection Curves

Ensemble of 5 pre-trained ResNets fine-tuned on Food101. Upper aleatoric uncertainty is used as rejection criterion.



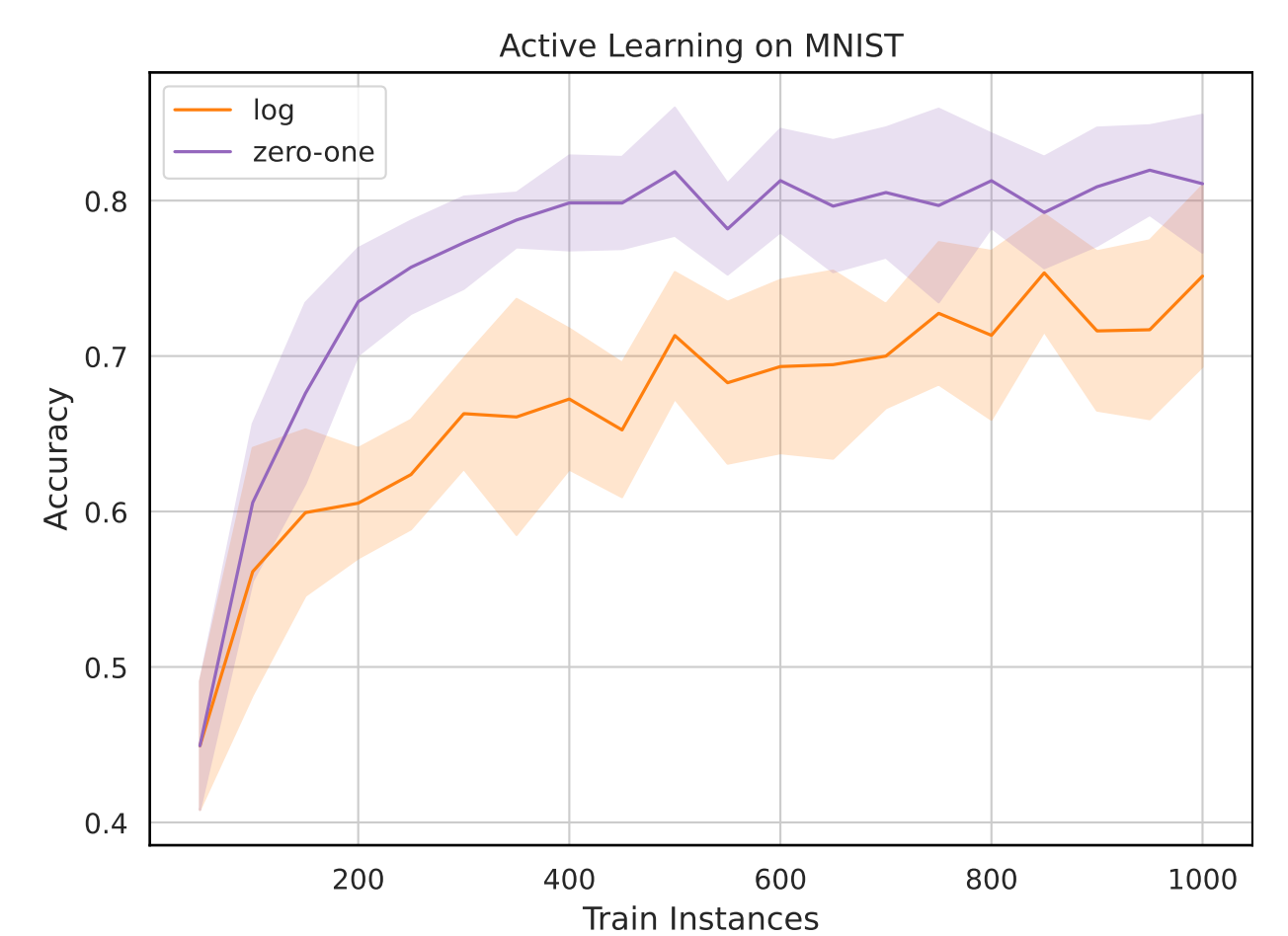
Out-of-Distribution Detection

Ensemble of 5 LeNets for FMNIST and 5 ResNets for Food101. OoD using epistemic uncertainty.

iD	OoD	log	Brier	spherical	zero-one	Hartley	entropy
FMNIST	MNIST	0.871 ± 0.017	0.812 ± 0.019	0.818 ± 0.019	0.742 ± 0.015	0.815 ± 0.019	0.826 ± 0.019
	KMNIST	0.973 ± 0.002	0.94 ± 0.004	0.946 ± 0.003	0.863 ± 0.006	0.944 ± 0.004	0.942 ± 0.003
Food101	SVHN	0.700 ± 0.04	0.572 ± 0.027	0.588 ± 0.038	0.669 ± 0.007	0.479 ± 0.023	0.681 ± 0.07
	CIFAR-100	0.805 ± 0.015	0.66 ± 0.016	0.681 ± 0.016	0.697 ± 0.008	0.576 ± 0.022	0.775 ± 0.021

Active Learning

Ensemble of 10 small MLPs. Initial training pool is 50 instances. Every round 50 new instances are acquired based on their epistemic uncertainty.



Future Work

- Can we **offer guarantees** on the credal set [5]?
- How does relaxing the **convexity assumption** affect the theoretical and empirical results?

