

# Is Epistemic Uncertainty Faithfully Represented by Evidential Deep Learning Methods?

WUML 2024

---

Mira Jürgens<sup>1</sup>, Nis Meinert<sup>2</sup>, Viktor Bengs<sup>3</sup>, Eyke Hüllermeier<sup>3</sup> and Willem Waegeman<sup>1</sup>

<sup>1</sup>Gent University <sup>2</sup>German Aerospace Center (DLR) <sup>3</sup>LMU Munich



Deutsches Zentrum  
für Luft- und Raumfahrt



# Setting

| Notation                                                                                                                   | Meaning                                                         |
|----------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------|
| $\mathcal{X} = \mathbb{R}^D$                                                                                               | instance space                                                  |
| $\mathcal{Y} = \begin{cases} \mathbb{R}^{D'}, & \text{regression} \\ \{1, \dots, K\}, & \text{classification} \end{cases}$ | label space                                                     |
| $\mathcal{D}_N = \{(\mathbf{x}_i, y_i)\}_{i=1}^N$                                                                          | training data                                                   |
| $P$                                                                                                                        | joint prob. distribution on $\mathcal{X} \times \mathcal{Y}$    |
| $p(\cdot \mathbf{x})$                                                                                                      | cond. density for $\mathbf{x} \in \mathcal{X}$ on $\mathcal{Y}$ |
| $\mathbb{P}(A)$                                                                                                            | set of all prob. distributions on $A$                           |

# What is Evidential Deep Learning

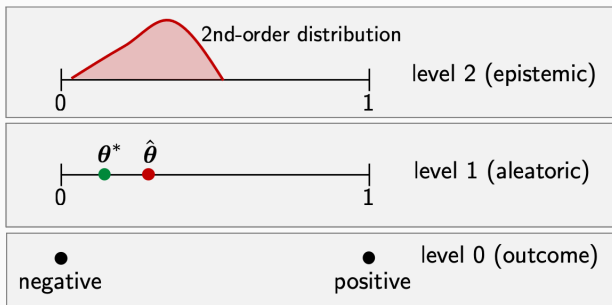
---

# Evidential Deep Learning for Classification

**Idea of EDL:** Learn a second-order distribution directly by which both aleatoric and epistemic uncertainty can be predicted.

## Dirichlet-Categorical Model

- level 1:  $y \sim \text{Cat}(\theta)$  with  $\theta \in \Delta_K$
- level 2:  $\theta \sim \text{Dir}(\mathbf{m})$  with  $\mathbf{m} \in \mathbb{R}_+^K$



**Figure 1:** Different levels of predictions<sup>1</sup>, exemplary for the binary classification case.

<sup>1</sup>Wimmer et al., *Quantifying Aleatoric and Epistemic Uncertainty in Machine Learning: Are Conditional Entropy and Mutual Information Appropriate Measures?*

# Evidential Deep Learning for Regression

## Normal-Normal-Inverse-Gamma Model

- level 1:  $y \sim \mathcal{N}(\mu, \sigma^2)$  with  $\mu \in \mathbb{R}$ ,  $\sigma^2 \in \mathbb{R}_+$
- level 2:  $(\mu, \sigma^2) \sim N\Gamma^{-1}(\mathbf{m})$  with  $\mathbf{m} = (\gamma, \nu, \alpha, \beta) \in \mathbb{R} \times \mathbb{R}_+^3$

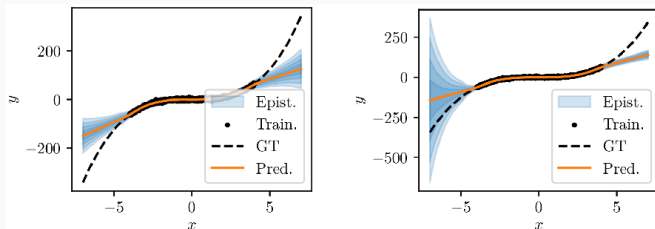


Figure 2: Results of Deep Evidential Regression<sup>2</sup> reproduced for two different runs<sup>3</sup>.

<sup>2</sup>Amini et al., “Deep evidential regression”.

<sup>3</sup>Meinert, Gawlikowski, and Lavin, “The Unreasonable Effectiveness of Deep Evidential Regression”.

# Different Problems, Different Distributions

For the following sections, let

- $\theta \equiv \theta(\mathbf{x}, \phi) \in \Theta$ , with  $\theta : \mathcal{X} \rightarrow \Theta$  be the parameters of the first-order predictor
- $\mathbf{m} \equiv \mathbf{m}(\mathbf{x}, \phi) \in \mathcal{M}$  with  $\mathbf{m} : \mathcal{X} \rightarrow \mathcal{M}$  the parameters of the second-order model

| problem                           | likelihood                                                                                               | conjugate prior                                                                                                            |
|-----------------------------------|----------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------|
| classification <sup>4</sup>       | $y \sim \text{Cat}(\theta)$<br>$\theta \in \Delta_k$                                                     | $\theta \sim \text{Dir}(\mathbf{m})$<br>$\mathbf{m} \in \mathbb{R}_+^K$                                                    |
| regression (univ.) <sup>5</sup>   | $y \sim \mathcal{N}(\mu, \sigma^2)$<br>$\mu \in \mathbb{R}, \sigma^2 \in \mathbb{R}_+$                   | $\theta \sim N\Gamma^{-1}(\mathbf{m})$<br>$\mathbf{m} = (\gamma, \nu, \alpha, \beta) \in \mathbb{R} \times \mathbb{R}_+^3$ |
| regression (multiv.) <sup>6</sup> | $\mathbf{y} \sim \mathcal{N}(\mu, \Sigma)$<br>$\mu \in \mathbb{R}^D, \Sigma \in \mathbb{R}^{D \times D}$ | $\theta \sim \text{NIW}(\mathbf{m})$<br>$\mathbf{m} = (\mu_0, \kappa, \nu, \Phi)$                                          |
| point processes                   | $y \sim \text{Pois}(\theta)$<br>$\theta \in \mathbb{R}_+$                                                | $\theta \sim \Gamma(\mathbf{m})$<br>$\mathbf{m} = (\alpha, \beta) \in \mathbb{R}_+^2$                                      |

<sup>4</sup>Sensoy, Kaplan, and Kandemir, “Evidential Deep Learning to Quantify Classification Uncertainty”.

<sup>5</sup>Amini et al., “Deep evidential regression”.

<sup>6</sup>Meinert and Lavin, “Multivariate deep evidential regression”.

While showing good results in downstream tasks, several critical analyses show theoretical flaws of EDL:

- Non-Convergence in the Classification case<sup>7</sup>
- Non-Convergence in the Regression Case<sup>8</sup>
- Non-Properness of its Loss functions<sup>9</sup>

—→ **Our approach:** Does the learned second-order distribution represent the (epistemic) uncertainty of the first-order parameters in a faithful way?

---

<sup>7</sup>Bengs, Hüllermeier, and Waegeman, “Pitfalls of Epistemic Uncertainty Quantification through Loss Minimisation”.

<sup>8</sup>Meinert, Gawlikowski, and Lavin, “The Unreasonable Effectiveness of Deep Evidential Regression”.

<sup>9</sup>Bengs, Hüllermeier, and Waegeman, *On Second-Order Scoring Rules for Epistemic Uncertainty Quantification*.

## **Comparing First and Second-Order Risk Minimization**

---

# First-Order vs Second-Order Risk Minimization

## First-Order Risk Minimization

- Hypothesis space

$$\mathcal{H}_1 \subset \mathbb{P}(\mathcal{Y}) = \{h_1 : \Theta \rightarrow \mathbb{P}(\mathcal{Y})\}$$

- Loss function (e.g. NLL)

$$L_1 : \mathcal{Y} \times \mathbb{P}(\mathcal{Y}) \rightarrow \mathbb{R}$$

- **Risk minimization** over parameters  $\Phi$

$$\min_{\Phi} \sum_{i=1}^N L_1(y_i, p(y | \theta(x_i; \Phi)))$$

# First-Order vs Second-Order Risk Minimization

## First-Order Risk Minimization

- Hypothesis space

$$\mathcal{H}_1 \subset \mathbb{P}(\mathcal{Y}) = \{h_1 : \Theta \rightarrow \mathbb{P}(\mathcal{Y})\}$$

- Loss function (e.g. NLL)

$$L_1 : \mathcal{Y} \times \mathbb{P}(\mathcal{Y}) \rightarrow \mathbb{R}$$

- **Risk minimization** over parameters  $\Phi$

$$\min_{\Phi} \sum_{i=1}^N L_1(y_i, p(y | \theta(x_i; \Phi)))$$

## Second-Order Risk Minimization

- Hypothesis space

$$\mathcal{H}_2 \subset \mathbb{P}(\Theta) = \{h_2 : \mathcal{M} \rightarrow \mathbb{P}(\Theta)\}$$

- Loss function (e.g. NLL)

$$L_1 : \mathcal{Y} \times \mathbb{P}(\mathcal{Y}) \rightarrow \mathbb{R}$$

- **Inner expectation minimization**

$$\min_{\Phi} \sum_{i=1}^N L_1 \left( y_i, \mathbb{E}_{\theta \sim p(\theta | m(x_i; \Phi))} [p(y | \theta)] \right) + \lambda R(\Phi).$$

- **Outer expectation minimisation**

$$\min_{\Phi} \sum_{i=1}^N \mathbb{E}_{\theta \sim p(\theta | m(x_i; \Phi))} \left[ L_1(y_i, p(y | \theta)) \right] + \lambda R(\Phi).$$

# Regularization

**Objective:** Enforce  $p(\theta | \mathbf{m}(\mathbf{x}_i; \Phi))$  to look similar to a predefined distribution  $p(\theta | \mathbf{m}_0)$  by minimising per-instance KL-divergence:

$$R(\Phi) = \sum_{i=1}^N d_{\text{KL}}(p(\theta | \mathbf{m}(\mathbf{x}_i; \Phi)), p(\theta | \mathbf{m}_0)).$$

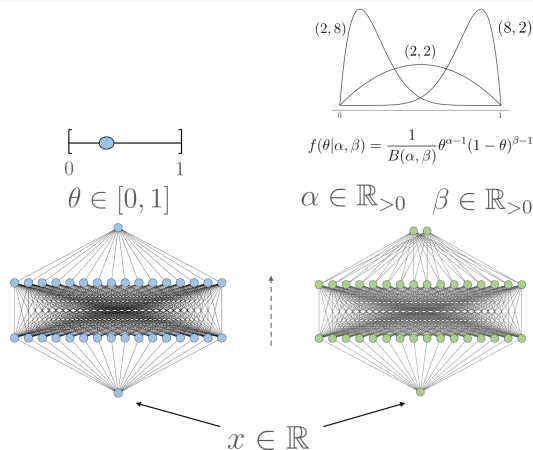
- choose  $p(\theta | \mathbf{m}_0)$  to parametrise uniform distribution (on first-order distributions)
- minimizing **KL-divergence**  $\leftrightarrow$  maximizing **entropy**
- other regularization losses exist, e.g. evidential loss<sup>10</sup>

$$R(\Phi) = \sum_{i=1}^N |y_i - \mathbb{E}[\mu_i]| \Psi = \sum_{i=1}^N |y_i - \gamma_i| (2\nu + \alpha)$$

---

<sup>10</sup>Amini et al., “Deep evidential regression”.

# Example: First vs Second Order Risk Minimization



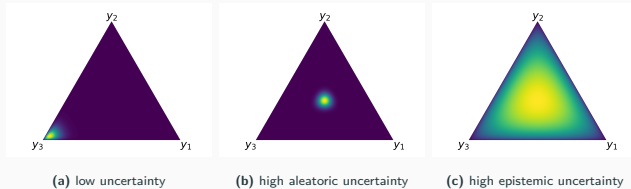
**Figure 3:** Overview of an exemplary neural network architecture for first (left) and second (right) order risk minimization for the case of binary classification.

**Do the resulting second-order distributions represent the underlying epistemic uncertainty of the first-order predictor?**

---

# What is a Faithful Representation of Epistemic Uncertainty?

Spread of the distribution should yield a valid estimate of EU:



**Figure 4:** Second-order predictions for the classification case with  $K = 3$ . See also Posterior network<sup>11</sup>, evidential DNN<sup>12</sup>

<sup>11</sup>Charpentier, Zügner, and Günnemann, “Posterior network: Uncertainty estimation without OOD samples via density-based pseudo-counts”.

<sup>12</sup>Sensoy, Kaplan, and Kandemir, “Evidential Deep Learning to Quantify Classification Uncertainty”.

<sup>13</sup>Bengs, Hüllermeier, and Waegeman, “Pitfalls of Epistemic Uncertainty Quantification through Loss Minimisation”.

# What is a Faithful Representation of Epistemic Uncertainty?

Spread of the distribution should yield a valid estimate of EU:

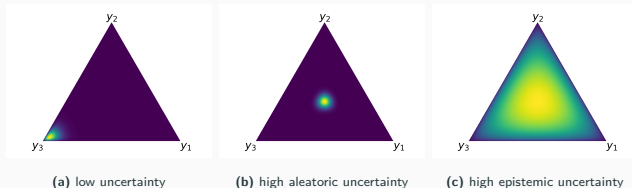


Figure 4: Second-order predictions for the classification case with  $K = 3$ . See also Posterior network<sup>11</sup>, evidential DNN<sup>12</sup>

Desirable convergence properties<sup>13</sup>:

1. Monotonicity: decreasing uncertainty with increasing sample size  $N$
2. Convergence to Dirac delta distribution when  $N \rightarrow \infty$

<sup>11</sup>Charpentier, Zügner, and Günnemann, “Posterior network: Uncertainty estimation without OOD samples via density-based pseudo-counts”.

<sup>12</sup>Sensoy, Kaplan, and Kandemir, “Evidential Deep Learning to Quantify Classification Uncertainty”.

<sup>13</sup>Bengs, Hüllermeier, and Waegeman, “Pitfalls of Epistemic Uncertainty Quantification through Loss Minimisation”.

# What is a Faithful Representation of Epistemic Uncertainty?

How can we directly evaluate the faithfulness of the resulting second-order distribution?

# What is a Faithful Representation of Epistemic Uncertainty?

How can we directly evaluate the faithfulness of the resulting second-order distribution?

## Def. 3.1: Reference Distribution

Let  $\theta_{\mathcal{D}_N}(\mathbf{x}; \Phi_{\mathcal{D}_N})$  denote the minimizer of the first order minimization problem for a training set  $\mathcal{D}_N$  of size  $N$ , where  $\mathcal{D}_N \sim P^N$ . For a given  $\mathbf{x} \in \mathcal{X}$ , we define the *reference second-order distribution* as

$$q_N(\theta | \mathbf{x}) := \int_{(\mathcal{X} \times \mathcal{Y})^N} \mathbb{I}[\theta_{\mathcal{D}_N}(\mathbf{x}; \Phi_{\mathcal{D}_N}) = \theta] dP^N, \quad (1)$$

where  $\mathbb{I}[\cdot]$  is the indicator function.

# What is a Faithful Representation of Epistemic Uncertainty?

How can we directly evaluate the faithfulness of the resulting second-order distribution?

## Def. 3.1: Reference Distribution

Let  $\theta_{\mathcal{D}_N}(\mathbf{x}; \Phi_{\mathcal{D}_N})$  denote the minimizer of the first order minimization problem for a training set  $\mathcal{D}_N$  of size  $N$ , where  $\mathcal{D}_N \sim P^N$ . For a given  $\mathbf{x} \in \mathcal{X}$ , we define the *reference second-order distribution* as

$$q_N(\theta | \mathbf{x}) := \int_{(\mathcal{X} \times \mathcal{Y})^N} \mathbb{I}[\theta_{\mathcal{D}_N}(\mathbf{x}; \Phi_{\mathcal{D}_N}) = \theta] dP^N, \quad (1)$$

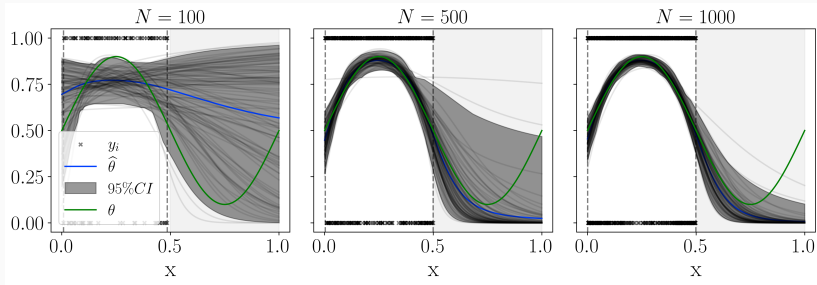
where  $\mathbb{I}[\cdot]$  is the indicator function.

$\implies$  **Can be approximated empirically** by resampling datasets

$\mathcal{D}_k = \{(x_{i,k}, y_{i,k})\}_{i=1}^N$  for  $k = 1, \dots, d \in \mathbb{N}$ , resulting in risk-minimizers  $(\hat{\theta}_1, \dots, \hat{\theta}_d)$  and

$$\hat{Q}_{N,d}(\theta | \mathbf{x}) := \frac{1}{d} \sum_{i=1}^d \mathbb{I}_{\hat{\theta}_i(\mathbf{x}, \phi) \leq \theta}$$

# Reference Distribution: Example



**Figure 5:** Example of empirically obtaining the reference distribution of the parameter  $\theta \in [0, 1]$  of a Bernoulli distribution, where  $\theta = 0.5 \cdot (0.8 \cdot \sin(2\pi x) + 1)$ . Empirical 95% confidence bounds are obtained by resampling the training data 100 times.

# Theoretical Results

---

# Theoretical Results 1

Let  $\mathcal{H}_J$  be the joint hypothesis space of  $\mathcal{H}_1$  and  $\mathcal{H}_2$  containing operators  $\mathcal{M} \rightarrow \mathbb{P}(\mathcal{Y})$  defined by

$$p(y \mid \mathbf{m}(\mathbf{x})) = \mathbb{E}_{\theta \sim p(\theta \mid \mathbf{m}(\mathbf{x}; \Phi))} p(y \mid \theta).$$

## Theorem 3.2: Unregularized Inner Loss Minimisation

Let  $\mathcal{C}_1$  and  $\mathcal{C}_J$  be the co-domains of  $\mathcal{H}_1$  and  $\mathcal{H}_J$ . The following properties hold when  $\mathcal{H}_2$  consists of

- Dirichlet distributions:  $\mathcal{H}_J$  is **not injective**, and  $\mathcal{C}_1 = \mathcal{C}_J$ ;
- NIG distributions:  $\mathcal{H}_J$  is **not injective**, and  $\mathcal{C}_1 \subset \mathcal{C}_J$ ;
- Gamma distributions:  $\mathcal{H}_J$  is **injective**, and  $\mathcal{C}_1 \subset \mathcal{C}_J$ .

# Theoretical Results 1

Let  $\mathcal{H}_J$  be the joint hypothesis space of  $\mathcal{H}_1$  and  $\mathcal{H}_2$  containing operators  $\mathcal{M} \rightarrow \mathbb{P}(\mathcal{Y})$  defined by

$$p(y \mid \mathbf{m}(\mathbf{x})) = \mathbb{E}_{\theta \sim p(\theta \mid \mathbf{m}(\mathbf{x}; \Phi))} p(y \mid \theta).$$

## Theorem 3.2: Unregularized Inner Loss Minimisation

Let  $\mathcal{C}_1$  and  $\mathcal{C}_J$  be the co-domains of  $\mathcal{H}_1$  and  $\mathcal{H}_J$ . The following properties hold when  $\mathcal{H}_2$  consists of

- Dirichlet distributions:  $\mathcal{H}_J$  is **not injective**, and  $\mathcal{C}_1 = \mathcal{C}_J$ ;
- NIG distributions:  $\mathcal{H}_J$  is **not injective**, and  $\mathcal{C}_1 \subset \mathcal{C}_J$ ;
- Gamma distributions:  $\mathcal{H}_J$  is **injective**, and  $\mathcal{C}_1 \subset \mathcal{C}_J$ .

**Implications: non-identifiability** for Dirichlet, NIG priors, leading to a wide spectrum of possible values for entropy of  $p(\theta \mid \mathbf{m}(\mathbf{x}_i, \phi))$

## Theoretical Results 2

Generalization of Theorem 1 of Bengs et al<sup>14</sup> for all exponential family members:

### Theorem 3.3: Unregularized Outer Expectation Minimisation

*Let  $L_1 : \mathcal{Y} \times \mathbb{P}(\mathcal{Y})$  be convex in its second argument, and let  $\mathcal{H}_2$  be a hypothesis space that contains universal approximators. A parameter set  $\Phi$  that yields  $p(\theta | \mathbf{m}(\mathbf{x}_i; \hat{\Phi})) = \delta(\theta(\mathbf{x}_i))$  for all  $i \in \{1, \dots, N\}$  is a minimizer of the outer expectation minimization problem, where  $\delta$  is the **Dirac delta function** on  $\theta(\mathbf{x}_i) \in \Theta$ .*

---

<sup>14</sup>Bengs, Hüllermeier, and Waegeman, “Pitfalls of Epistemic Uncertainty Quantification through Loss Minimisation”.

Generalization of Theorem 1 of Bengs et al<sup>14</sup> for all exponential family members:

### Theorem 3.3: Unregularized Outer Expectation Minimisation

*Let  $L_1 : \mathcal{Y} \times \mathbb{P}(\mathcal{Y})$  be convex in its second argument, and let  $\mathcal{H}_2$  be a hypothesis space that contains universal approximators. A parameter set  $\Phi$  that yields  $p(\theta | \mathbf{m}(\mathbf{x}_i; \hat{\Phi})) = \delta(\theta(\mathbf{x}_i))$  for all  $i \in \{1, \dots, N\}$  is a minimizer of the outer expectation minimization problem, where  $\delta$  is the **Dirac delta function** on  $\theta(\mathbf{x}_i) \in \Theta$ .*

**Implications:** one of the risk minimizers is a Dirac distribution..

---

<sup>14</sup>Bengs, Hüllermeier, and Waegeman, “Pitfalls of Epistemic Uncertainty Quantification through Loss Minimisation”.

### Theorem 3.3: Regularized Loss Minimisation

*Consider the optimization problems for inner and outer expectation minimization with the entropy regularization of the second-order distribution. Then, for any data-generating distribution  $P$ , there exists  $\lambda \geq 0$  and  $\mathbf{x} \in \mathcal{X}$  for which the second-order distribution differs from the reference distribution in Def 3.1.*

### Theorem 3.3: Regularized Loss Minimisation

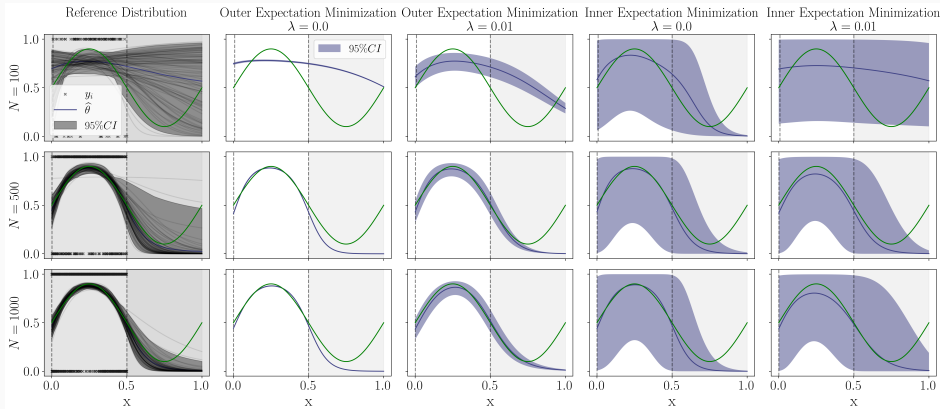
*Consider the optimization problems for inner and outer expectation minimization with the entropy regularization of the second-order distribution. Then, for any data-generating distribution  $P$ , there exists  $\lambda \geq 0$  and  $\mathbf{x} \in \mathcal{X}$  for which the second-order distribution differs from the reference distribution in Def 3.1.*

**Implications:**  $\lambda \in \mathbb{R}_{\geq 0}$  defines an *uncertainty budget* that cannot be exceeded for the given amount of training data points!

## Empirical Results

---

# Experimental Results: Classification



**Figure 6:** Binary classification experiments for training sample size  $N \in \{100, 500, 1000\}$ , using synthetic training data  $\mathcal{D}_N = \{(x_i, y_i)\}_{i=1}^N$ , where  $x_i \sim U([0, 0.5])$ ,  $y_i \sim \text{Bern}(\theta(x_i))$ . The underlying Bernoulli parameter is defined by  $\theta(x_i) = 0.5 \cdot (0.8 \cdot \sin(2\pi x_i) + 1)$ .

→ more experimental results in

<https://arxiv.org/pdf/2402.09056.pdf>

## Conclusion

---

## Summary

- DER does not result in distributions that are faithfully representing EU
- The regularization parameter  $\lambda$  yields an "uncertainty budget" that cannot be exceeded for a given amount of training data points
- While outer loss minimisation can lead to a convergence to a Dirac delta distribution, inner loss minimisation leads to training instability

## Future outlook:

- further analysis is needed, especially for different regularization losses
- distance-based approaches to quantify distance to reference distribution
- experimental analysis of heteroscedastic case

# Conclusion

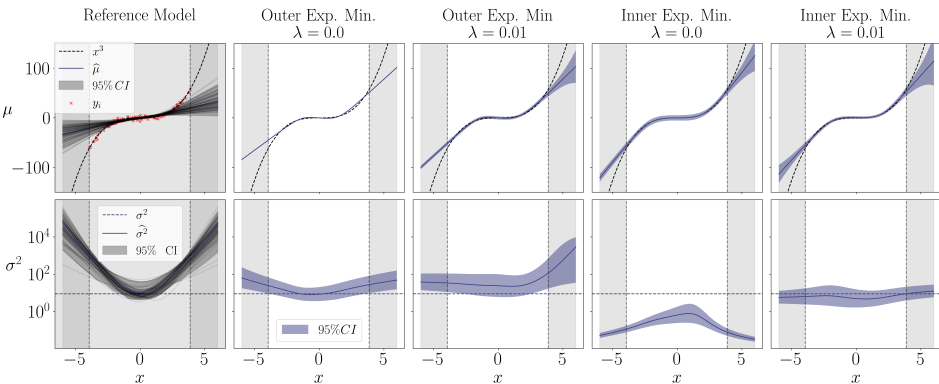
For regression results, proofs, more theoretical analysis, check out our paper:



THANK YOU!

# Empirical Results: Regression

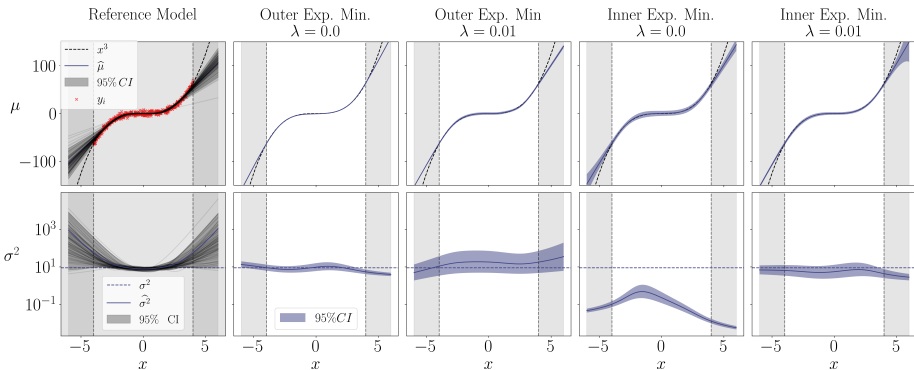
$N = 100$ :



**Figure 7:** Regression experiment with  $\mathcal{D}_N = \{(x_i, x_i^3 + \epsilon)\}_{i=1}^N$  for  $N \in \{100, 500, 1000\}$ , where  $x_i \in U([-4, 4])$ ,  $\epsilon \sim \mathcal{N}(0, \sigma^2 = 9)$ . The reference model learns the parameters  $\theta = (\mu, \sigma)$  of the underlying normal distribution, by optimizing the negative log-likelihood.

# Empirical Results: Regression

$N = 500$ :



# Empirical Results: Regression

$N = 1000$ :

