

Quantification of Prediction Uncertainty using Random Fuzzy Sets

Thierry Denœux

Université de technologie de Compiègne, France

Zhejiang University of Technology, Hangzhou, July 16, 2026

Need for uncertainty quantification

Prediction uncertainty

- **Uncertainty quantification** is a fundamental problem in statistics and AI.
- In prediction tasks such as, e.g., weather forecasting, we need to tell the user how **confident** they can be in the predictions.

Aleatory uncertainty

- Inherent variability
- Cannot be reduced
- Usually represented by probability

Epistemic uncertainty

- Lack of knowledge
- Can be reduced by information
- Imprecision or incompleteness of information

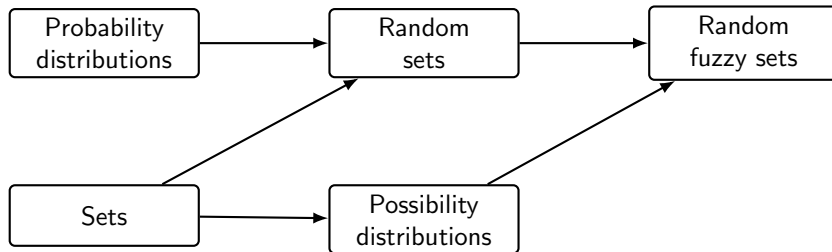
Why probability is not be enough

- Probability distributions are natural for representing random uncertainties.
- But epistemic uncertainties are very important in statistical and model-based prediction, because of
 - ▶ Limited sample size,
 - ▶ Poor data quality,
 - ▶ Model misspecification, etc.
- A single probability distribution may then give a false impression of precision.

Random fuzzy sets

Random fuzzy sets are designed to represent both aleatory and epistemic uncertainties simultaneously.

A hierarchy of uncertainty models



- Random sets extend probability and sets by allowing set-valued realisations.
- Possibility distributions encode graded compatibility.
- Random fuzzy sets combine both ideas.

Outline

- 1 Uncertainty representation frameworks
 - From mass functions to random sets
 - From random sets to random fuzzy sets
 - Gaussian random fuzzy models
- 2 Application to machine learning
 - Distance-based approach
 - Likelihood-based approach

Outline

- 1 Uncertainty representation frameworks
 - From mass functions to random sets
 - From random sets to random fuzzy sets
 - Gaussian random fuzzy models
- 2 Application to machine learning

Classical Dempster–Shafer theory

- Let Y be an unknown quantity with value in a finite frame Θ .
- Evidence about Y can be represented by a **mass function** assigning weights to subsets of Θ :

$$m(A) \geq 0, \quad \sum_{A \subseteq \Theta} m(A) = 1.$$

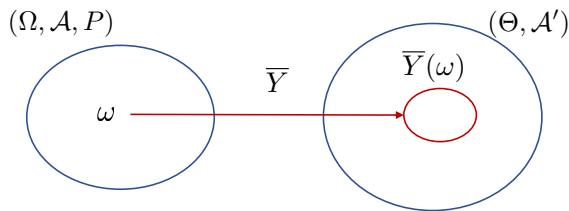
- $m(A)$ is the probability that the evidence only says that $Y \in A$.
- **Belief** and **plausibility** are lower and upper probabilities:

$$\text{Bel}_m(B) = \sum_{A \subseteq B} m(A), \quad \text{Pl}_m(B) = \sum_{A \cap B \neq \emptyset} m(A) = 1 - \text{Bel}_m(B^c).$$

- Interpretation:
 - ▶ $\text{Bel}_m(B)$: degree of total support for B ;
 - ▶ $\text{Bel}_m(B)$: degree of compatibility of B with the evidence.

From finite frames to infinite spaces

Many unknowns take values in $\mathbb{R}^p$, not in a finite set.



- The mass function is replaced by the more general concept of random subset of Θ .
- This leads to the notion of a **random set**, defined by a probability space $(\Omega, \mathcal{A}, P)$ and a set-valued function

$$\overline{Y} : \Omega \rightarrow \mathcal{P}(\Theta).$$

Interpretation

- Here, Ω is a set of **interpretations** of the evidence.
- For each elementary interpretation ω , we observe a set-valued constraint:

$$Y \in \overline{Y}(\omega).$$

Belief and plausibility for a random set

For $B \subseteq \Theta$,

$$\text{Bel}_{\bar{Y}}(B) = P\{\bar{Y} \subseteq B\}, \quad \text{Pl}_{\bar{Y}}(B) = P\{\bar{Y} \cap B \neq \emptyset\} = 1 - \text{Bel}_{\bar{Y}}(B^c).$$

Belief

Probability that the evidence **supports** B .

Plausibility

Probability that the evidence is **compatible** with B (does not support B^c).

Properties:

- $\text{Bel}(\emptyset) = \text{Pl}(\emptyset) = 0$
- $\text{Bel}(\Theta) = \text{Pl}(\Theta) = 1$
- $\forall B \subseteq \Theta, \text{Bel}(B) \leq \text{Pl}(B)$.
- If the focal sets are singletons, $\text{Bel} = \text{Pl}$ is additive.

Random sets as evidence on infinite frames

Key equivalence

A random set is a probability distribution over crisp focal sets.

- This extends belief functions from finite frames to general spaces.
- But it still assumes that each focal element is crisp, i.e., the evidence imposes hard constraints on the unknown quantity.

What if the evidence is itself vague or graded, such as, e.g.,
“ Y is close to 1.5”, or “ Y is small”?

From crisp focal sets to fuzzy focal sets

Random set

$$\overline{Y}(\omega) \subseteq \Theta$$

A crisp constraint:

$$Y \in \overline{Y}(\omega).$$

Random fuzzy set

$$\tilde{Y}(\omega) : \Theta \rightarrow [0, 1]$$

A soft (graded, flexible) constraint:

$$Y = y \text{ is possible to degree } \tilde{Y}(\omega)(y).$$

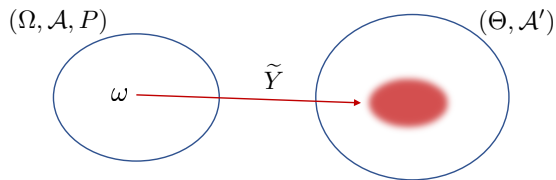
Possibility and necessity generated by a fuzzy set

- A **fuzzy subset** of Θ is a map $\pi : \Theta \rightarrow [0, 1]$. It is **normal** if $\sup_{\theta \in \Theta} \pi(\theta) = 1$.
- A normal fuzzy set representing a soft constraint on an unknown X is called a **possibility distribution**.
- A possibility distribution induces **possibility and necessity measures**:

$$\Pi_{\pi}(B) = \sup_{x \in B} \pi(x), \quad N_{\pi}(B) = 1 - \Pi_{\pi}(B^c).$$

- ▶ $\Pi_{\pi}(B)$ measures the extent to which B is **compatible** with the constraint.
- ▶ $N_{\pi}(B)$ measures the extent to which B is **implied** by the constraint.
- A fuzzy focal set $\tilde{Y}(\omega)$ therefore induces possibility and necessity measures, conditional on ω .

Definition of a random fuzzy set



Random fuzzy set

A random fuzzy set is defined by a probability space $(\Omega, \mathcal{A}, P)$ and a fuzzy-valued mapping

$$\tilde{Y} : \Omega \rightarrow \mathcal{F}(\Theta),$$

where $\tilde{Y}(\omega)$ is a fuzzy subset of Θ .

- Random sets are recovered when all fuzzy focal sets are crisp.
- Random variables are recovered when all fuzzy focal sets are points (singletons).
- Possibility distributions are recovered when there is only one non-random focal fuzzy set.

Belief and plausibility induced by a RFS

- For any $\omega \in \Omega$, $\tilde{Y}(\omega)$ is a possibility distribution for X .
- For any $B \subseteq \Theta$, under measurability conditions, $\Pi_{\tilde{Y}(\omega)}(B)$ and $N_{\tilde{Y}(\omega)}(B)$ are random variables.
- We define:

$$\text{Pl}_{\tilde{Y}}(B) = \mathbb{E} [\Pi_{\tilde{Y}}(B)] ,$$

$$\text{Bel}_{\tilde{Y}}(B) = \mathbb{E} [N_{\tilde{Y}}(B)] = 1 - \text{Pl}_{\tilde{Y}}(B^c).$$

Interpretation

- Plausibility is the average possibility of B over all uncertain interpretations.
- Belief is the average necessity of B .

Contour function

The plausibility of a singleton gives the **contour function**:

$$\text{pl}_{\tilde{\gamma}}(y) = \text{Pl}_{\tilde{\gamma}}(\{y\}) = \mathbb{E} \left[\tilde{Y}(y) \right].$$

- It is the average membership degree of x over the random fuzzy focal sets.
- It generalises both probability density-like information and possibility distributions.
- It is often easier to visualise than the full random fuzzy set.

One object, several limiting cases

Model	RFS representation	Main meaning
Random variable	random singleton	precise randomness
Random set	random crisp set	imprecise evidence
Possibility distribution	fixed fuzzy set	vague evidence
Random fuzzy set	random fuzzy set	variability + imprecision + vagueness

Unifying view

RFSs extend Dempster–Shafer theory in the same way that fuzzy sets extend crisp sets: focal sets become graded.

Combination of random fuzzy sets

- Let $\tilde{Y}_1$ and $\tilde{Y}_2$ be two independent random fuzzy sets on Θ .
- Their combination $\tilde{Y}_1 \oplus \tilde{Y}_2$ is obtained by intersecting their fuzzy focal sets:

$$\tilde{Y}_1 \oplus \tilde{Y}_2 : (\omega_1, \omega_2) \mapsto \tilde{Y}_1(\omega_1) \cap \tilde{Y}_2(\omega_2).$$

- Intersection is defined as the normalised product:

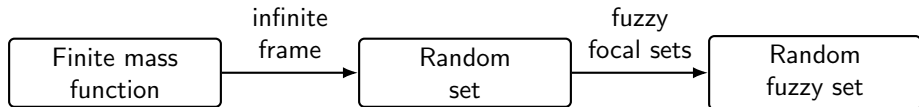
$$(\tilde{Y}_1(\omega_1) \cap \tilde{Y}_2(\omega_2))(x) = \frac{\tilde{Y}_1(\omega_1)(x) \tilde{Y}_2(\omega_2)(x)}{\sup_{u \in \Theta} \tilde{Y}_1(\omega_1)(u) \tilde{Y}_2(\omega_2)(u)}.$$

- The product probability measure $P_1 \otimes P_2$ is updated to account for the conflict.

Main idea

$\oplus$ combines two independent random fuzzy sets by intersecting their focal fuzzy sets, and renormalising.

Conceptual summary



- Random set: probability distribution over sets.
- Random fuzzy set: probability distribution over possibility distributions.
- Belief and plausibility remain lower and upper probabilities.

Need for practical continuous models

- General RFSs are very flexible but difficult to specify and propagate.
- For numerical uncertainty propagation, we need parametric models.
- **Gaussian random fuzzy numbers and vectors** provide such a tractable family.
- They are indexed by three parameters representing
 - ▶ location,
 - ▶ dispersion,
 - ▶ imprecision.

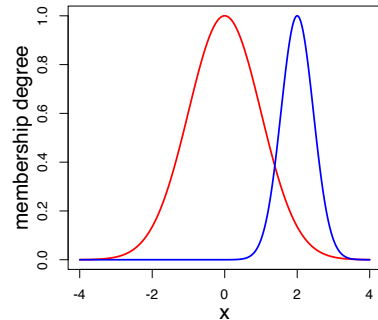
Gaussian fuzzy number

Definition

A **Gaussian fuzzy number (GFN)** is defined by

$$\tilde{y}(y) = \exp\left(-\frac{h}{2}(y - m)^2\right),$$

where $m \in \mathbb{R}$ is the **mode** and $h \geq 0$ is a **precision** parameter.



We write

$$\tilde{y} = \text{GFN}(m, h).$$

Gaussian random fuzzy number

Definition

A **Gaussian random fuzzy number (GRFN)** is a GFN whose mode is a normal random variable

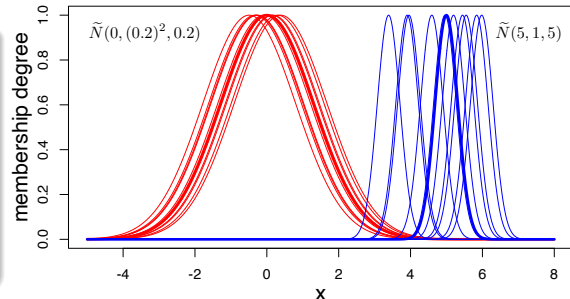
$$\tilde{Y}(M) = \text{GFN}(M, h),$$

with

$$M \sim \mathcal{N}(\mu, \sigma^2), \quad h \geq 0.$$

We write

$$\tilde{Y} \sim \tilde{\mathcal{N}}(\mu, \sigma^2, h).$$



Interpreting the parameters

- μ is the central value.
- σ^2 measures how much the modes varies across interpretations.
- h measures the specificity of each focal fuzzy number.
- Three limiting cases:

Condition	Interpretation
$\sigma^2 = 0$	Gaussian fuzzy number
$h \rightarrow \infty$	Gaussian random variable
$h = 0$	vacuous information

Message

A GRFN continuously interpolates between probability, possibility and ignorance.

Gaussian fuzzy vector

Definition

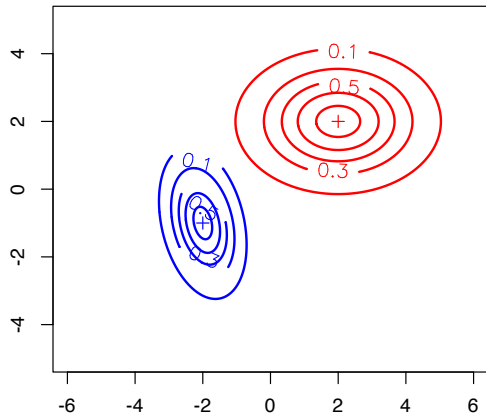
A **Gaussian fuzzy vector (GFV)** is defined by

$$\tilde{\mathbf{y}}(\mathbf{y}) = \exp\left(-\frac{1}{2}(\mathbf{y} - \mathbf{m})^T \mathbf{H}(\mathbf{y} - \mathbf{m})\right),$$

where $\mathbf{m} \in \mathbb{R}^p$ is the mode and $\mathbf{H} \in \mathbb{R}^{p \times p}$ is a positive definite **precision matrix**.

We write

$$\tilde{\mathbf{y}} = \text{GFV}(\mathbf{m}, \mathbf{H}).$$



Gaussian random fuzzy vector

Definition

A **Gaussian random fuzzy vector (GRFV)** is a GFV whose mode is a normal random vector:

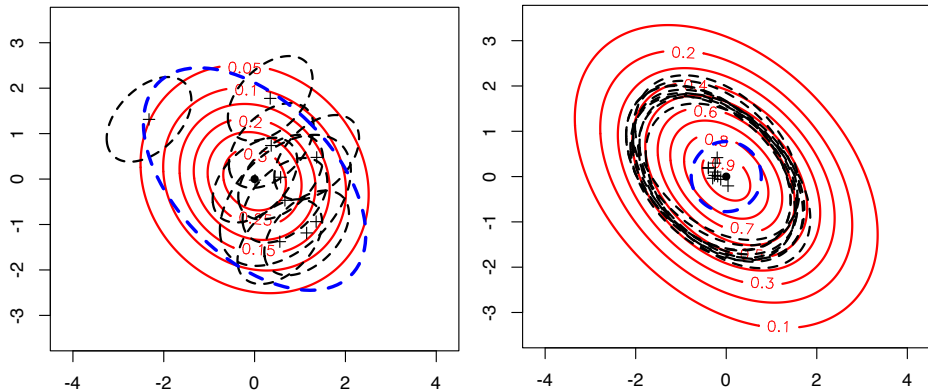
$$\tilde{\mathbf{Y}}(\mathbf{M}) = \text{GFV}(\mathbf{M}, \mathbf{H}), \quad \mathbf{M} \sim \mathcal{N}(\boldsymbol{\mu}, \boldsymbol{\Sigma}), \quad \mathbf{H}, \boldsymbol{\Sigma} \succeq 0,$$

We write

$$\tilde{\mathbf{Y}} \sim \tilde{\mathcal{N}}(\boldsymbol{\mu}, \boldsymbol{\Sigma}, \mathbf{H}).$$

- $\boldsymbol{\Sigma}$ controls the dispersion of the random centres.
- $\mathbf{H}$ controls the shape and precision of each fuzzy focal ellipsoid.

Representation of GRFVs

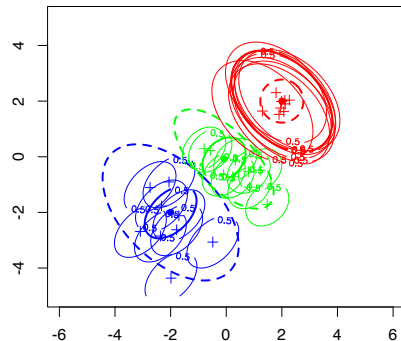
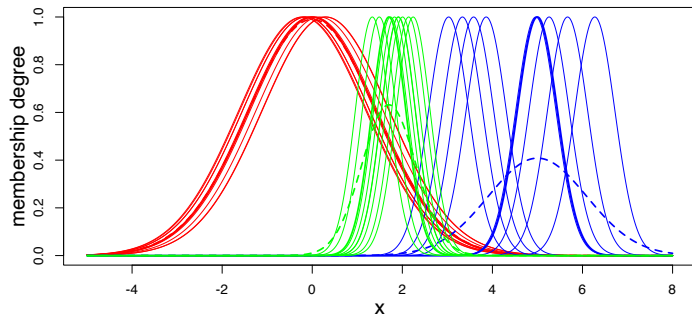


For each GRFV, we show 0.5-level curves of ten focal sets (black broken curves), a contour plot of the contour function (red solid curves), as well as a 95% confidence ellipse for the mode $\mathbf{M} \sim N(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ (blue broken curves).

Combination of Gaussian RFSs

Closeness property

GRFNs and GRFVs are closed under product-intersection.



Interpretability

- From $\tilde{Y}(\mathbf{x}) \sim \tilde{\mathcal{N}}\{\mu(\mathbf{x}), \sigma^2(\mathbf{x}), h(\mathbf{x})\}$, we can extract:
 - ▶ a point prediction $\mu(\mathbf{x})$;
 - ▶ a contour function $\text{pl}_{\mathbf{x}}(y)$;
 - ▶ lower and upper CDFs;
 - ▶ belief intervals at prescribed confidence levels;
 - ▶ a measure of epistemic imprecision through $h(\mathbf{x})$.
- These summaries are interpretable by users who do not know the underlying theory.

Outline

1 Uncertainty representation frameworks

2 Application to machine learning

- Distance-based approach
- Likelihood-based approach

Machine learning: the setting

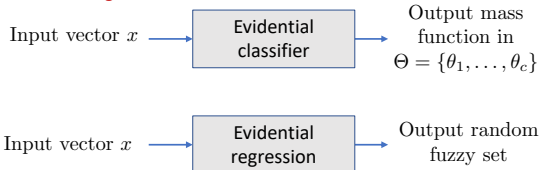
- Machine learning builds predictive models from data for classification, regression, decision support and knowledge discovery.
- In **supervised learning**, the learning set is

$$\mathcal{L} = \{(\mathbf{x}_i, y_i), i = 1, \dots, n\}.$$

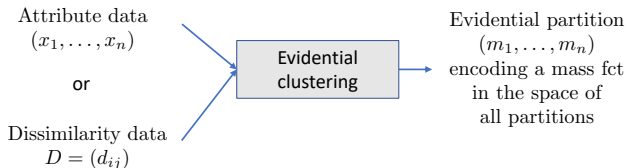
- The response may be
 - ▶ a class label: $Y \in \Theta = \{\theta_1, \dots, \theta_c\}$;
 - ▶ a numerical value: $Y \in \Theta \subseteq \mathbb{R}$.
- Evidence theory enriches prediction by returning not only a label or value, but also a **belief function** describing predictive uncertainty.

Evidential machine learning

Supervised learning



Unsupervised learning



Distance-based evidential prediction

- The starting idea is simple:

nearby training points or prototypes provide stronger evidence.

- A prototype close to the query input $\mathbf{x}$ gives an informative predictive statement.
- A prototype far from $\mathbf{x}$ gives weak evidence or even vacuous information.
- This produces a natural warning in regions with little or no training data.

Principle

Distances to prototypes induce pieces of evidence; these pieces are combined into a predictive belief function.

Prototype evidence as a GRFN

- The leaning vectors are summarised by K prototypes $\mathbf{w}_1, \dots, \mathbf{w}_K$.
- Prototype k supplies

$$\tilde{Y}_k(\mathbf{x}) \sim \tilde{\mathcal{N}}\{\mu_k(\mathbf{x}), \sigma_k^2, s_k(\mathbf{x})h_k\}.$$

- A simple local mean is affine:

$$\mu_k(\mathbf{x}) = \beta_k^T \mathbf{x} + \beta_{k0}.$$

- The similarity term is

$$s_k(\mathbf{x}) = \exp\{-\gamma_k \|\mathbf{x} - \mathbf{w}_k\|^2\}.$$

- Close prototypes give high epistemic precision; distant prototypes give weak evidence.

Combining GRFN prototype outputs

A simplified product-intersection rule yields

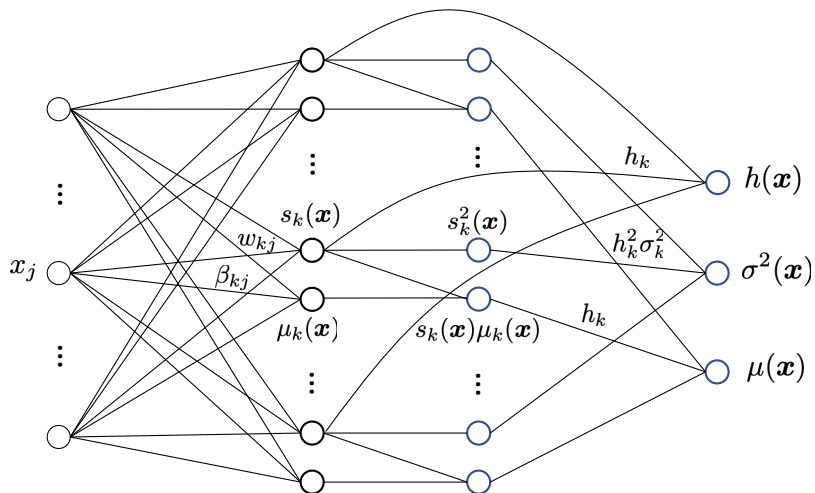
$$\tilde{Y}(\mathbf{x}) = \tilde{Y}_1(\mathbf{x}) \boxplus \dots \boxplus \tilde{Y}_K(\mathbf{x}) \sim \tilde{\mathcal{N}}(\mu(\mathbf{x}), \sigma^2(\mathbf{x}), h(\mathbf{x})),$$

with

$$\mu(\mathbf{x}) = \frac{\sum_{k=1}^K s_k(\mathbf{x}) h_k \mu_k(\mathbf{x})}{\sum_{k=1}^K s_k(\mathbf{x}) h_k}, \quad h(\mathbf{x}) = \sum_{k=1}^K s_k(\mathbf{x}) h_k,$$
$$\sigma^2(\mathbf{x}) = \frac{\sum_{k=1}^K s_k(\mathbf{x})^2 h_k^2 \sigma_k^2}{\left(\sum_{k=1}^K s_k(\mathbf{x}) h_k\right)^2}.$$

- $\sigma^2(\mathbf{x})$ captures aleatory uncertainty.
- $h(\mathbf{x})$ captures epistemic uncertainty.

ENNreg architecture



Learning

- The model is trained by minimising

$$C_{\lambda, \epsilon, \xi, \rho}^{(R)}(\theta) = \frac{1}{n} \sum_{i=1}^n \ell_{\lambda, \epsilon}(y_i, \tilde{Y}(\mathbf{x}_i; \theta)) + \frac{\xi}{K} \sum_k h_k + \frac{\rho}{K} \sum_k \gamma_k,$$

where

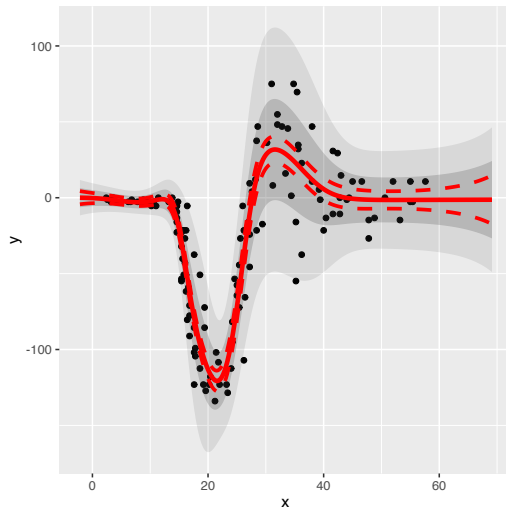
$$\ell_{\lambda, \epsilon} = -\lambda \log \text{Bel}_{\tilde{Y}(\mathbf{x}_i)}([y_i]_{\epsilon}) - (1 - \lambda) \log \text{Pl}_{\tilde{Y}(\mathbf{x}_i)}([y_i]_{\epsilon}),$$

and $[y_i]_{\epsilon} = [y_i - \epsilon, y_i + \epsilon]$.

- Hyperparameters:
 - ▶ ϵ accounts for the finite precision of continuous observations (typically, $\epsilon = 0.01 \hat{\sigma}_Y$);
 - ▶ λ controls the trade-off between precise and cautious predictions (typically, $\lambda = 0.9$);
 - ▶ ξ and ρ regularise, respectively, the effective number of prototypes and the departure from a linear model. They are determined by cross-validation.

Example: motorcycle data

- Point predictions (solid red curve),
- lower and upper expectations (dashed red curves)
- belief intervals at levels 0.50 and 0.95 (shaded areas).



Comparison with classical methods (RMS)

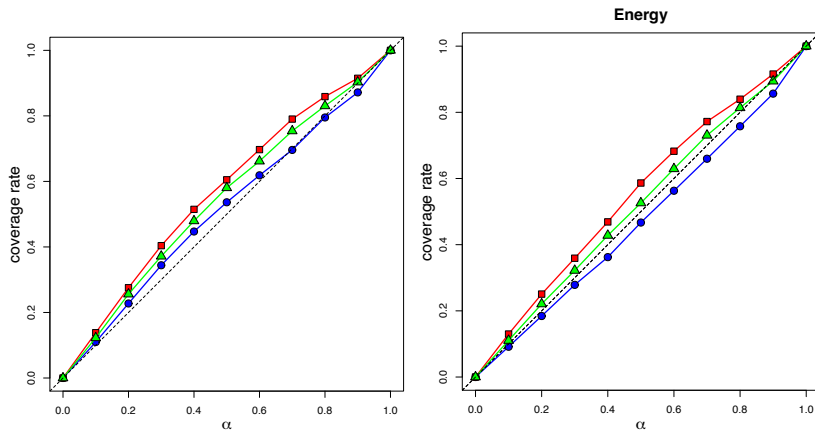
	ENNreg	RBF	RVM	SVM	GP	RF	MLP
Boston	2.87 ± 0.14	3.31 ± 0.19	3.42 ± 0.17	3.17 ± 0.15	3.70 ± 0.22	3.11 ± 0.14	3.14 ± 0.14
Energy	1.06 ± 0.05	2.06 ± 0.08	1.79 ± 0.05	1.39 ± 0.06	2.58 ± 0.07	1.75 ± 0.06	0.95 ± 0.16
Concr.	5.10 ± 0.12	6.30 ± 0.19	6.38 ± 0.16	5.62 ± 0.13	6.93 ± 0.13	4.64 ± 0.12	4.82 ± 0.16
Yacht	0.44 ± 0.04	2.00 ± 0.20	1.88 ± 0.20	1.93 ± 0.11	6.12 ± 0.31	0.96 ± 0.08	0.50 ± 0.05
Wine	0.63 ± 0.01	0.63 ± 0.01	0.80 ± 0.02	0.61 ± 0.01	0.61 ± 0.01	0.56 ± 0.01	0.77 ± 0.01
kin8nm	0.08 ± 0.00	0.11 ± 0.00	–	0.09 ± 0.00	0.08 ± 0.00	0.14 ± 0.00	0.07 ± 0.00
Crime	0.14 ± 0.00	0.14 ± 0.00	0.14 ± 0.00	0.14 ± 0.00	0.14 ± 0.00	0.14 ± 0.00	0.14 ± 0.00
Resid.	0.11 ± 0.01	0.16 ± 0.01	0.17 ± 0.01	0.15 ± 0.01	0.22 ± 0.01	0.16 ± 0.01	0.14 ± 0.01
Airfoil	1.46 ± 0.03	1.70 ± 0.04	2.58 ± 0.04	2.37 ± 0.04	2.49 ± 0.04	1.44 ± 0.04	1.53 ± 0.04
Bike	6.59 ± 0.19	6.49 ± 0.15	6.64 ± 0.14	7.11 ± 0.16	7.55 ± 0.14	6.86 ± 0.17	9.68 ± 0.20

Comparison with SOTA methods (RMS & NLL)

	RMS				
	ENNreg	PBP	MC-dropout	Deep ens.	Deep ev. reg.
Boston	2.87 $\pm$ 0.14	3.01 $\pm$ 0.18	2.97 $\pm$ 0.19	3.28 $\pm$ 1.00	3.06 $\pm$ 0.16
Energy	1.06 $\pm$ 0.05	1.80 $\pm$ 0.05	1.66 $\pm$ 0.04	2.09 $\pm$ 0.29	2.06 $\pm$ 0.10
Concr.	5.10 $\pm$ 0.12	5.67 $\pm$ 0.09	5.23 $\pm$ 0.12	6.03 $\pm$ 0.58	5.85 $\pm$ 0.15
Yacht	0.44 $\pm$ 0.04	1.02 $\pm$ 0.05	1.11 $\pm$ 0.09	1.58 $\pm$ 0.48	1.57 $\pm$ 0.56
Wine	0.63 $\pm$ 0.01	0.64 $\pm$ 0.01	0.62 $\pm$ 0.01	0.64 $\pm$ 0.04	0.61 $\pm$ 0.02
kin8nm	0.08 $\pm$ 0.00	0.10 $\pm$ 0.00	0.10 $\pm$ 0.00	0.09 $\pm$ 0.00	0.09 $\pm$ 0.00

	NLL				
	ENNreg	PBP	MC-dropout	Deep ens.	Deep ev. reg.
Boston	2.53 $\pm$ 0.07	2.57 $\pm$ 0.09	2.46 $\pm$ 0.06	2.41 $\pm$ 0.25	2.35 $\pm$ 0.06
Energy	1.14 $\pm$ 0.07	2.04 $\pm$ 0.02	1.99 $\pm$ 0.02	1.38 $\pm$ 0.22	1.39 $\pm$ 0.06
Concr.	3.38 $\pm$ 0.13	3.16 $\pm$ 0.02	3.04 $\pm$ 0.02	3.06 $\pm$ 0.18	3.01 $\pm$ 0.02
Yacht	0.13 $\pm$ 0.12	1.63 $\pm$ 0.02	1.55 $\pm$ 0.03	1.18 $\pm$ 0.21	1.03 $\pm$ 0.19
Wine	0.94 $\pm$ 0.01	0.97 $\pm$ 0.01	0.93 $\pm$ 0.01	0.94 $\pm$ 0.12	0.89 $\pm$ 0.05
kin8nm	-1.19 $\pm$ 0.00	-0.90 $\pm$ 0.01	-0.95 $\pm$ 0.01	-1.20 $\pm$ 0.02	-1.24 $\pm$ 0.01

Calibration plots



Probabilistic predictions (blue), raw evidential predictions (red) and adjusted evidential predictions (green).

A different starting point

- The distance-based approach is non-parametric.
- The **likelihood-based approach** starts from a statistical model, making it possible to derive a **prediction equation**

$$Y = \varphi(\mathbf{x}, \boldsymbol{\theta}, U),$$

where U is a noise variable and $\boldsymbol{\theta}$ is unknown.

- The data inform us about $\boldsymbol{\theta}$ through the likelihood.
- Prediction is obtained by propagating parameter uncertainty and noise through the prediction equation.

Likelihood and relative likelihood

Let $\mathbf{Y}$ have density f_{θ} and let $\mathbf{y}$ be observed.

Likelihood

$$L_{\mathbf{y}}(\theta) = \alpha f_{\theta}(\mathbf{y}),$$

where $\alpha > 0$ is arbitrary.

Relative likelihood

$$\tilde{\pi}_{\mathbf{y}}(\theta) = \frac{L_{\mathbf{y}}(\theta)}{\sup_{\theta' \in \Theta} L_{\mathbf{y}}(\theta')}.$$

- It is normalised: $\sup_{\theta} \tilde{\pi}_{\mathbf{y}}(\theta) = 1$.
- It ranks parameter values according to their compatibility with the data.

Possibilistic interpretation

- The relative likelihood is interpreted as a **possibility distribution** on the parameter space:

$$\tilde{\pi}_{\mathbf{y}} : \Theta \rightarrow [0, 1].$$

- $\tilde{\pi}_{\mathbf{y}}(\boldsymbol{\theta}) = 1$ for the most likely parameters.
- Small values mean that $\boldsymbol{\theta}$ is poorly supported by the data.
- This is a form of **Bayesian inference without a probabilistic prior**: if a probabilistic prior is added and combined by product-intersection, the Bayesian posterior is recovered.

Gaussian approximation

Let $\hat{\theta}$ be a maximiser of the log-likelihood. A second-order Taylor expansion gives

$$\log \tilde{\pi}_{\mathbf{y}}(\boldsymbol{\theta}) \approx -\frac{1}{2}(\boldsymbol{\theta} - \hat{\boldsymbol{\theta}})^T \mathbf{J}(\hat{\boldsymbol{\theta}})(\boldsymbol{\theta} - \hat{\boldsymbol{\theta}}),$$

where

$$\mathbf{J}(\hat{\boldsymbol{\theta}}) = - \left. \frac{\partial^2 \log \tilde{\pi}_{\mathbf{y}}}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^T} \right|_{\boldsymbol{\theta} = \hat{\boldsymbol{\theta}}}$$

is the **observed information matrix**.

Approximate possibilistic information on $\boldsymbol{\theta}$

$$\tilde{\pi}_{\mathbf{y}} \approx \text{GFV} \left(\hat{\boldsymbol{\theta}}, \mathbf{J}(\hat{\boldsymbol{\theta}}) \right).$$

Simple example: normal mean

- Suppose

$Y_1, \dots, Y_n$ i.i.d. from $\mathcal{N}(\theta, \sigma^2)$, σ^2 known.

- The relative likelihood is

$$\tilde{\pi}_{\mathbf{y}}(\theta) = \exp \left\{ -\frac{n}{2\sigma^2} (\theta - \bar{y})^2 \right\}.$$

- Hence the parameter uncertainty is represented exactly by a Gaussian possibility distribution centred at $\bar{y}$, with precision n/σ^2 .
- As n increases, epistemic imprecision decreases.

Prediction equation

- For a new input $\mathbf{x}$, write the predictive model as

$$Y = \varphi(\mathbf{x}, \boldsymbol{\theta}, U),$$

where U is random with known distribution.

- Parameter uncertainty is possibilistic:

$$\boldsymbol{\theta} \sim \tilde{\pi}_{\mathbf{y}}.$$

- Noise uncertainty is probabilistic:

$$U \sim P_U.$$

- Combining them naturally yields a random fuzzy set for Y .

Constructing the predictive RFS

- For each simulated noise value u , propagate the parameter possibility distribution:

$$\tilde{\pi}_{\mathbf{y},u}(y) = \sup_{\boldsymbol{\theta}: y=\varphi(\mathbf{x},\boldsymbol{\theta},u)} \tilde{\pi}_{\mathbf{y}}(\boldsymbol{\theta}).$$

- Since U is random, the map

$$\tilde{Y} : u \mapsto \tilde{\pi}_{\mathbf{y},u}$$

is a random fuzzy set.

- The resulting predictive RFS combines:
 - ▶ aleatory uncertainty from U ;
 - ▶ epistemic uncertainty from $\tilde{\pi}_{\mathbf{y}}$.

Linear-Gaussian case

- Suppose the prediction equation is linear:

$$Y = a(\mathbf{x}) + \mathbf{b}(\mathbf{x})^T \boldsymbol{\theta} + \sigma(\mathbf{x})U.$$

- If $U \sim N(0, 1)$ and $\tilde{\pi}_{\boldsymbol{\theta}}$ is Gaussian, the predictive RFS is a GRFN.
- This gives a closed-form approximation

$$\tilde{Y}(\mathbf{x}) \sim \tilde{\mathcal{N}}(\mu(\mathbf{x}), \sigma^2(\mathbf{x}), h(\mathbf{x})).$$

- The three parameters have direct predictive interpretations.
- These results hold approximately if the prediction equation is linearised about $\hat{\boldsymbol{\theta}}$.

Exact and approximate workflows

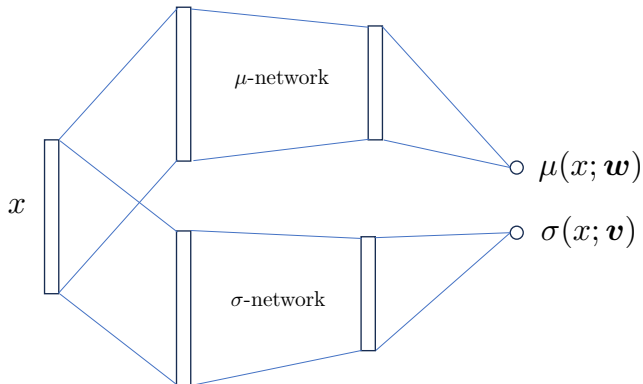
- Exact propagation is possible in simple models.
- Monte Carlo propagation can be used when closed forms are unavailable.
- Gaussian approximation is often sufficient for large samples or weakly nonlinear models.
- A practical workflow is:
 - ① fit the statistical model;
 - ② compute or approximate the relative likelihood;
 - ③ propagate through the prediction equation;
 - ④ summarise the predictive RFS by contours, CDF bounds, and belief intervals.

Heteroscedastic regression model

- In many prediction problems, the noise variance depends on the input.
- A neural network can output both

$$\mu(\mathbf{x}; \mathbf{w}), \quad \sigma^2(\mathbf{x}; \mathbf{v}).$$

- The prediction equation is
$$Y = \mu(\mathbf{x}; \mathbf{w}) + \sigma(\mathbf{x}; \mathbf{v})U, \quad U \sim \mathcal{N}(0, 1).$$
- The likelihood is used to build a possibility distribution over all network parameters.



Parameter possibility distribution

- Let $\hat{\boldsymbol{\theta}} = (\hat{\mathbf{w}}^T, \hat{\mathbf{v}}^T)^T$ be the maximum likelihood estimate.
- The relative likelihood is approximated by

$$\tilde{\pi}_{\mathbf{y}}(\boldsymbol{\theta}) \approx \exp\left\{-\frac{1}{2}(\boldsymbol{\theta} - \hat{\boldsymbol{\theta}})^T \mathbf{J}(\hat{\boldsymbol{\theta}})(\boldsymbol{\theta} - \hat{\boldsymbol{\theta}})\right\}.$$

- The information matrix has a block structure separating mean and variance parameters.
- This approximation turns parameter uncertainty into a tractable Gaussian possibility distribution.

Linearisation for prediction

- For a fixed input $\mathbf{x}$, linearise both networks around $\hat{\boldsymbol{\theta}}$:

$$\mu(\mathbf{x}; \mathbf{w}) \approx \mu(\mathbf{x}, \hat{\mathbf{w}}) + \mathbf{g}_{\mu}(\mathbf{x})^T (\mathbf{w} - \hat{\mathbf{w}}),$$

$$\sigma(\mathbf{x}; \mathbf{v}) \approx \sigma(\mathbf{x}, \hat{\mathbf{v}}) + \mathbf{g}_{\sigma}(\mathbf{x})^T (\mathbf{v} - \hat{\mathbf{v}}).$$

- This yields an approximate predictive equation

$$Y \approx \begin{pmatrix} \mathbf{g}_{\mu}(\mathbf{x}) \\ U \mathbf{g}_{\sigma}(\mathbf{x}) \end{pmatrix}^T \boldsymbol{\theta} + \mu(\mathbf{x}; \hat{\mathbf{w}}) - \mathbf{g}_{\mu}(\mathbf{x})^T \hat{\mathbf{w}} + U (\sigma(\mathbf{x}; \hat{\mathbf{v}}) - \mathbf{g}_{\sigma}(\mathbf{x})^T \hat{\mathbf{v}}), \quad U \sim \mathcal{N}(0, 1)$$

- This is linear in the full parameter vector $\boldsymbol{\theta}$. Hence Gaussian possibilistic propagation can be applied.
- The resulting RFS can be further approximated by a fixed-precision GRFN.

Interpretation of the predictive GRFN

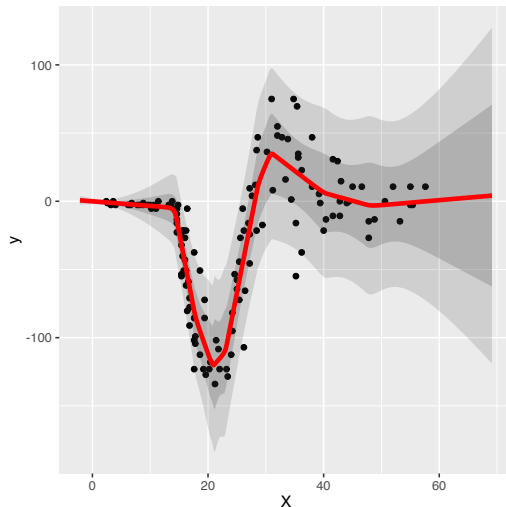
Final predictive approximation

$$\tilde{Y}(\mathbf{x}) \stackrel{\text{app.}}{\sim} \tilde{\mathcal{N}} \left(\mu(\mathbf{x}; \hat{\mathbf{w}}), \sigma^2(\mathbf{x}; \hat{\mathbf{v}}), h(\mathbf{x}, \hat{\boldsymbol{\theta}}) \right).$$

Component	Interpretation
$\mu(\mathbf{x}; \hat{\mathbf{w}})$	point prediction / expected response
$\sigma^2(\mathbf{x}; \hat{\mathbf{v}})$	aleatory variability of Y given x
$h(\mathbf{x}, \hat{\boldsymbol{\theta}})$	epistemic precision due to finite data and parameter uncertainty

Example: motorcycle data

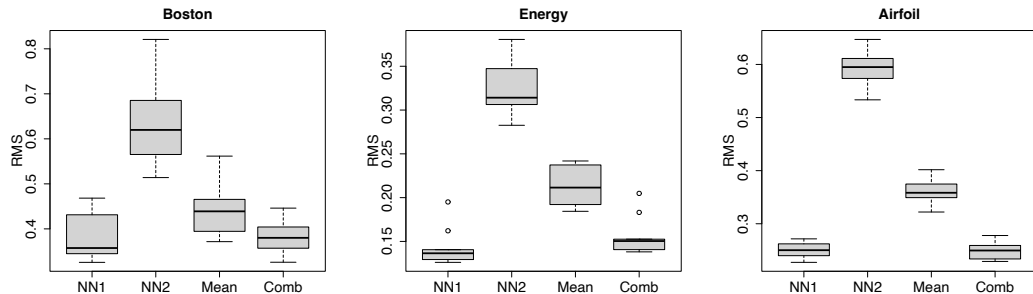
- Point predictions (solid red curve),
- belief intervals at levels 0.50 and 0.95 (shaded areas)
- μ -network: 20 hidden units;
 σ -network: 10 hidden units.



Experiment

- The reliability of predictions depends on the size of the training set.
- To simulate another fusion scenario involving the combination of predictions with different reliabilities, we combined the outputs of two neural networks, trained with 60% and 10% of the data.
- The process was repeated ten times with different random training/validation/test partitions.

Results



Box plots of root-mean-squared errors for the five datasets and the following methods: neural network trained on 60% of the data (NN1), neural network trained on 10% of the data (NN2), average prediction (Mean) and combination by the product-intersection rule (Comb).

Distance-based vs. likelihood-based approaches

Distance-based	Likelihood-based
Evidence induced by distance to prototypes	Evidence induced by relative likelihood of parameters
Good for flexible local models and deep features	Good when a prediction equation is available
Uncertainty learned from geometry and loss	Epistemic uncertainty tied to likelihood theory
Often easy to visualise	Often easier to justify statistically

Take-home messages

- Random fuzzy sets provide a language for predictive uncertainty with both aleatory and epistemic components.
- Belief and plausibility can be introduced operationally as lower and upper support for events.
- Distance-based methods build predictive RFSs from local evidence supplied by neighbours or learned prototypes.
- Likelihood-based methods build predictive RFSs by interpreting relative likelihood as a possibility distribution and propagating it through the prediction equation.
- For users, the practical outputs are point predictions, contours, lower/upper CDFs, belief intervals, and extrapolation warnings.

References

<https://www.hds.utc.fr/~tdenoeux>



T. Denœux

Reasoning with fuzzy and uncertain evidence using epistemic random fuzzy sets: general framework and practical models.

Fuzzy Sets and Systems, 453:1–36, 2023



T. Denœux

Quantifying Prediction Uncertainty in Regression using Random Fuzzy Sets: the ENNreg model

IEEE Transactions on Fuzzy Systems, 31(10):3690–3699, 2023



T. Denœux

Uncertainty Quantification in Regression Neural Networks using Evidential Likelihood-based Inference

International Journal of Approximate Reasoning, 182:109423, 2025