



Stony Brook University

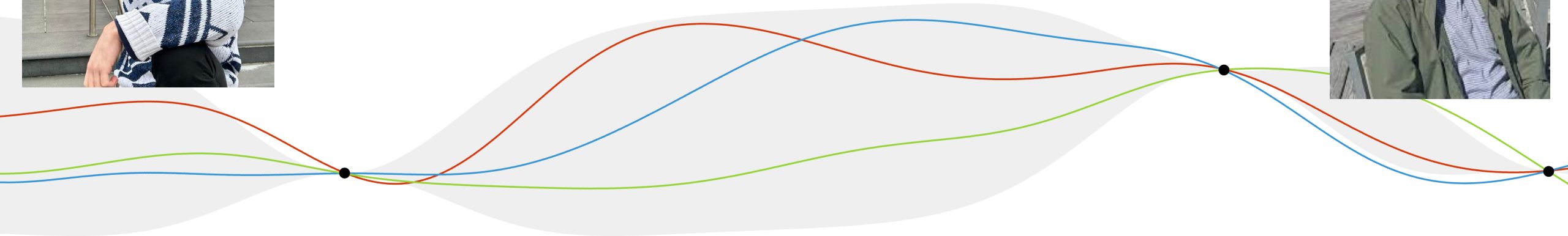


Explainable Learning with Gaussian Processes

Petar M. Djurić

Department of Electrical and Computer Engineering
Stony Brook University

Based on joint work with Kurt Butler and Guanchao Feng



Explainability in machine learning

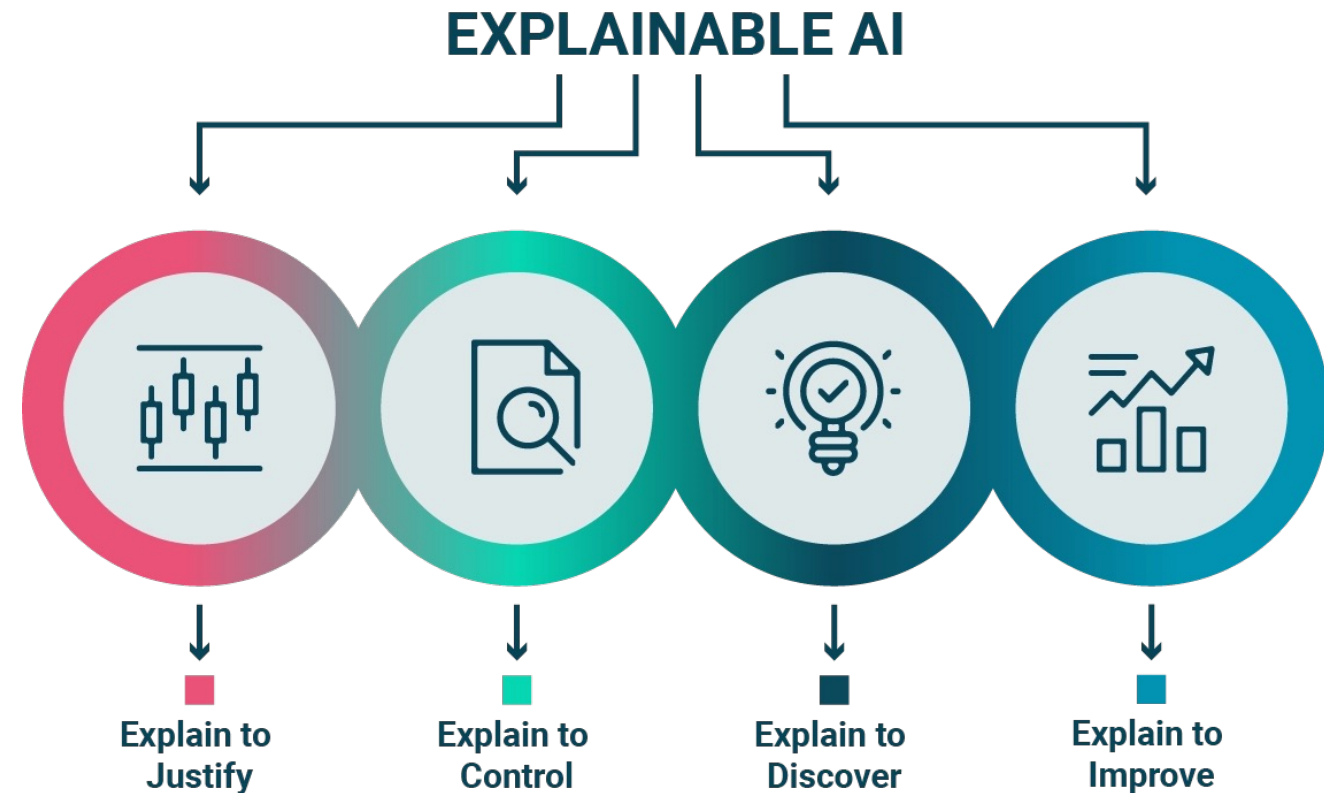
Explainability in machine learning refers to the ability to make a model's decisions or predictions understandable to humans.

Why it matters?

- Trust and transparency
- Debugging and improvement:
- Accountability and ethics:
- Fairness and bias detection:

Types of explainability:

- Global explainability
- Local explainability



The concept of attribution

Attribution refers to identifying the impact or contribution of input features on a model's output.

An example: *Medical diagnosis model*

Consider a model that predicts the likelihood of a disease based on **age, family history, blood pressure, and cholesterol levels**.

Suppose the model predicts a **70% probability** of the disease.

Attribution quantifies how much each feature contributed to the prediction. For instance:

- **High cholesterol** might contribute **30%**.
- **Age** could contribute **20%**.
- **Family history** could add **15%**.
- **Blood pressure** might add **5%**.



Image taken from the Internet.

Predicting house prices



Image taken from the Internet.

The model takes in features such as: **Square footage**, **Number of bedrooms**, **Location** (e.g., distance to city center), **Age of the house**, **Condition of the house**.

Now, suppose the model predicts that a specific house is worth **\$500,000**.

Attribution answers the question: *How much did each feature contribute to this price prediction?*

Loan approval

Suppose a machine learning model predicts a 70% chance of loan approval for a customer based on:

- Income
- Credit score
- Employment status

An attribution method provides values for the contribution of each feature, e.g.,

- Income: +20% (pushed the prediction up by 20%)
- Credit Score: +30%
- Employment Status: +20%



Image taken from Forbes.com

Approaches

- Gradient based methods (saliency maps, integrated gradients)
- Perturbation-based methods (SHAP - SHapley Additive exPlanations)
- Surrogate models (decision trees)
- Attention mechanisms (transformers in NLP)
- Layer-wise relevance propagation
- Model-specific methods (tree SHAP)

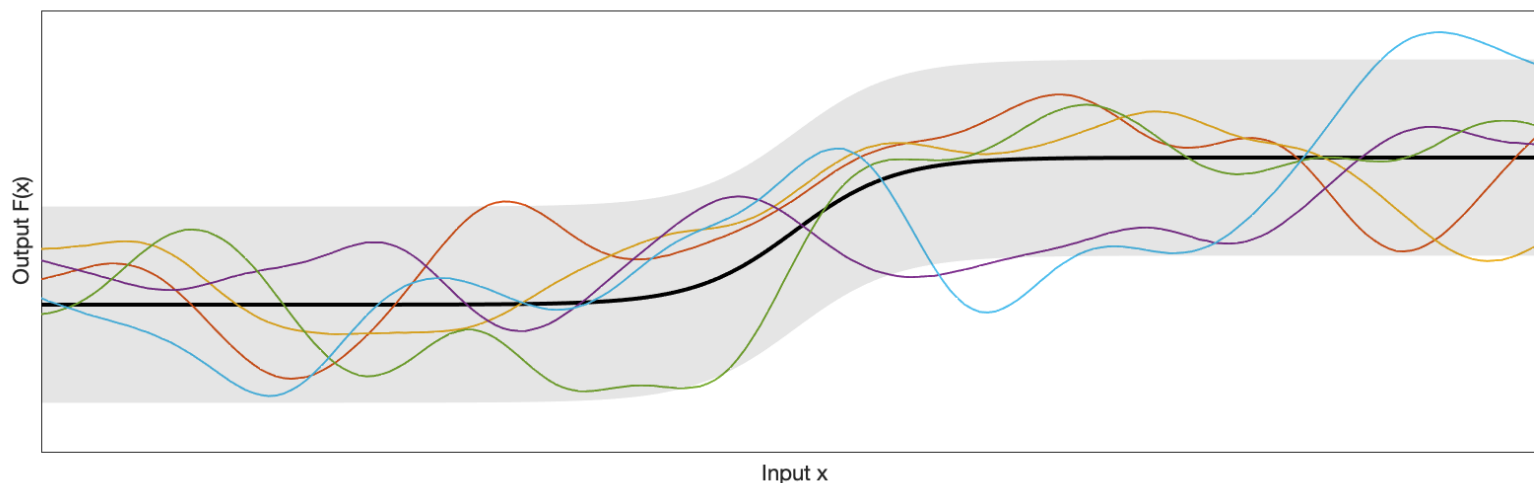


Lloyd Shapley
1923-1996

Our workhorse: Gaussian processes

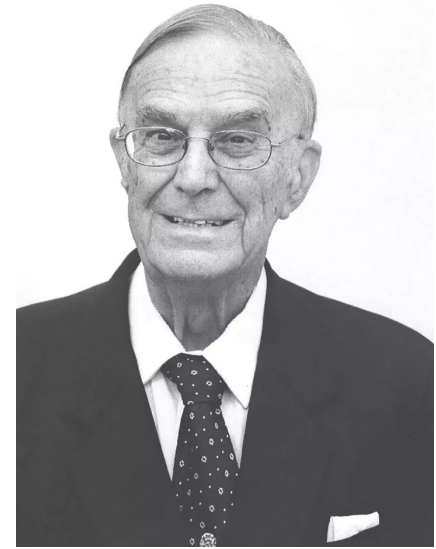
- A **Gaussian process** (GP) is a probability distribution over a space of functions.
 - A GP is specified by its mean and covariance functions, $m(\mathbf{x})$, $k(\mathbf{x}, \mathbf{x}')$.

$$F \sim \mathcal{GP}(m, k) \quad \longleftrightarrow \quad \begin{aligned} \mathbb{E}(F(\mathbf{x})) &= m(\mathbf{x}) \\ \text{Cov}(F(\mathbf{x}), F(\mathbf{x}')) &= k(\mathbf{x}, \mathbf{x}') \end{aligned}$$



A bit of history

Danie Gerhardus Krige (26 August 1919 – 3 March 2013) - a South African statistician and mining engineer; was professor at the University of the Witwatersrand, Republic of South Africa. In the 1950s, sought a more accurate method to estimate ore grades by considering the spatial structure and the variability of samples.

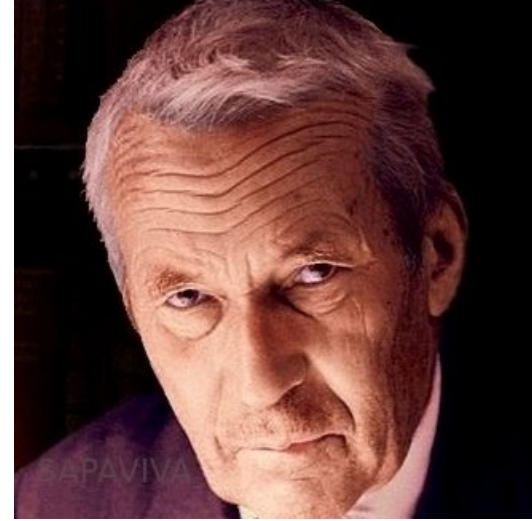


Georges François Paul Marie Matheron (2 December 1930 – 7 August 2000) was a French mathematician and civil engineer of mines. His PhD advisor was Paul Lévy. He established a rigorous mathematical framework for kriging, describing it as the **best linear unbiased estimator** for spatially correlated data. Published in **Traité de géostatistique appliquée**, Editions Technip 1962

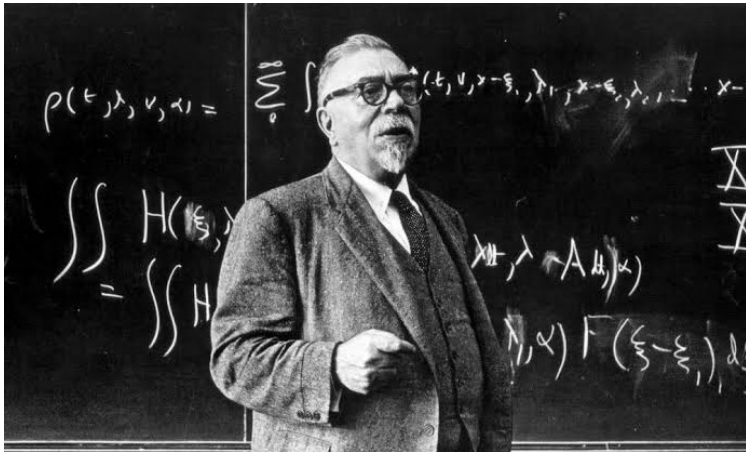
A bit more history



1777-1855



1903-1987



1894-1964



A recommendation for fn reading

the theory  
that would
not die  
how bayes' rule cracked
the enigma code,
hunted down russian
submarines & emerged
triumphant from two 
centuries of controversy
sharon bertsch mcgrayne

"If you're not thinking like a Bayesian, perhaps you should be."

—John Allen Paulos, New York Times Book Review



Sharon Bertsch McGrayne is the author of highly-praised books about scientific discoveries and the scientists who make them.

Gaussian process regression

- Gaussian process regression (GPR) is a Bayesian approach to learning an unknown functional relationship.

Model

$$y_n = F(\mathbf{x}_n) + \varepsilon_n$$
$$\varepsilon_n \sim \mathcal{N}(0, \sigma_n^2)$$

Likelihood

$$\mathbf{y} | F \sim \mathcal{N}(\mathbf{m}, \mathbf{K})$$

Goal: Obtain a posterior distribution over possible functions

$$\mathcal{D} = \{(\mathbf{x}_n, y_n) | n = 1, \dots, N\}$$

$$\mathbf{y} = \begin{bmatrix} y_1 \\ \vdots \\ y_N \end{bmatrix} \quad \mathbf{m} = \begin{bmatrix} m(\mathbf{x}_1) \\ \vdots \\ m(\mathbf{x}_N) \end{bmatrix}$$

$$\mathbf{K} = \begin{bmatrix} k(\mathbf{x}_1, \mathbf{x}_1) & \cdots & k(\mathbf{x}_1, \mathbf{x}_N) \\ \vdots & \ddots & \vdots \\ k(\mathbf{x}_N, \mathbf{x}_1) & \cdots & k(\mathbf{x}_N, \mathbf{x}_N) \end{bmatrix}$$

Gaussian Process Regression

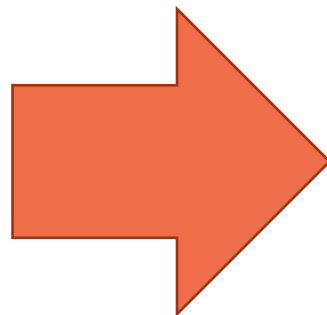
- Gaussian process regression (GPR) exploits GPs by recognizing that they are conjugate to Gaussian likelihoods.

Prior belief

$$F \sim \mathcal{GP}(m, k)$$

Likelihood

$$\mathbf{y}|F \sim \mathcal{N}(\mathbf{m}, \mathbf{K})$$



Posterior belief

$$F|\mathcal{D} \sim \mathcal{GP}(m_p, k_p)$$

$$m_p(\mathbf{x}) = m(\mathbf{x}) + \mathbf{k}(\mathbf{x})^\top (\mathbf{K} + \sigma_n^2 \mathbf{I})^{-1} (\mathbf{y} - \mathbf{m})$$

$$k_p(\mathbf{x}, \mathbf{x}') = k(\mathbf{x}, \mathbf{x}') + \mathbf{k}(\mathbf{x})^\top (\mathbf{K} + \sigma_n^2 \mathbf{I})^{-1} \mathbf{k}(\mathbf{x}')$$

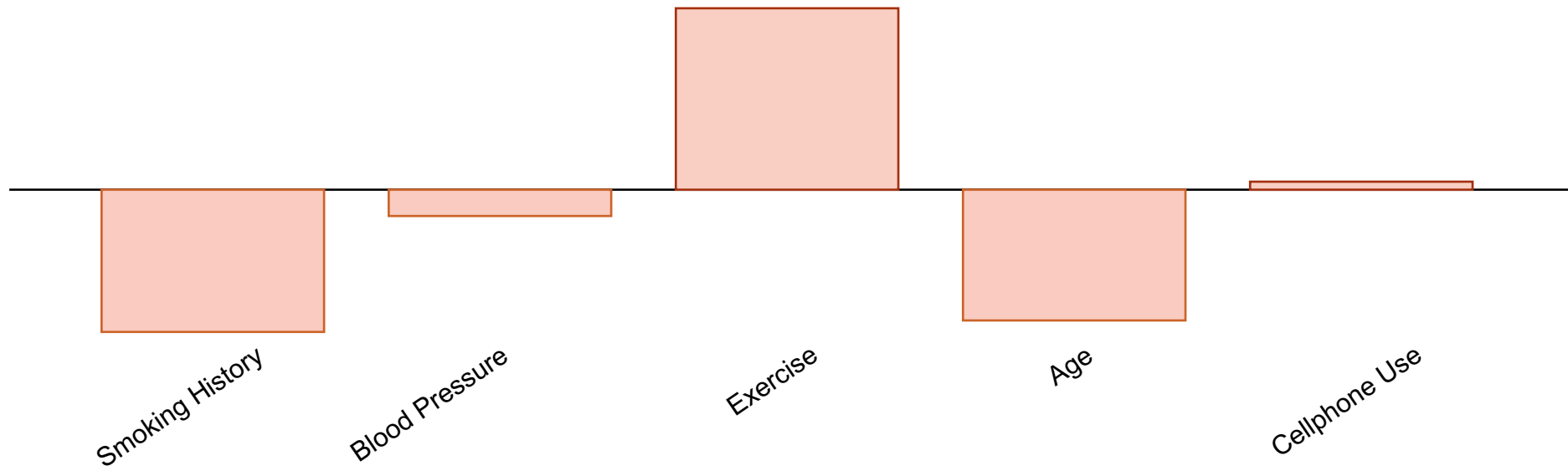
$$\mathbf{k}(\mathbf{x}) = \begin{bmatrix} k(\mathbf{x}, \mathbf{x}_1) \\ \vdots \\ k(\mathbf{x}, \mathbf{x}_N) \end{bmatrix}$$

Attribution Theory

- **Feature attributions** quantify how much each input variable contributed to the output result.

$$\underbrace{F(\mathbf{x})}_{\text{New prediction}} - \underbrace{F(\tilde{\mathbf{x}})}_{\text{Baseline point}} = \sum_{i=1}^D \underbrace{attr_i(\mathbf{x}|F)}_{\text{Contribution of feature } i}$$

Attributions to a
patient predicted
cancer risk



Attribution Theory

Consider the **Bayesian linear regression model**

$$y|\mathbf{x}, \mathbf{w} \sim \mathcal{N}(\mathbf{w}^\top \mathbf{x}, \sigma^2)$$

The predictions from the model are given by

$$F(\mathbf{x}) - F(\tilde{\mathbf{x}}) = \mathbf{w}^\top (\mathbf{x} - \tilde{\mathbf{x}}) = \sum_{i=1}^D w_i (x_i - \tilde{x}_i)$$

and the canonical choice for the attribution to the x_i feature is

$$attr_i(\mathbf{x}) = w_i (x_i - \tilde{x}_i)$$

Attribution Theory

Let

$$\mathbf{w} \sim \mathcal{N}(\boldsymbol{\mu}, \boldsymbol{\Sigma})$$

Then we can show that

$$\mathbf{w}|X, \mathbf{y} \sim \mathcal{N}(\boldsymbol{\mu}', \boldsymbol{\Sigma}')$$

where

$$\begin{aligned}\boldsymbol{\Sigma}' &= \left(\boldsymbol{\Sigma}^{-1} + \frac{1}{\sigma^2} X^\top X \right)^{-1} \\ \boldsymbol{\mu}' &= \boldsymbol{\Sigma}' \left(\boldsymbol{\Sigma}^{-1} \boldsymbol{\mu} + \frac{1}{\sigma^2} X^\top \mathbf{y} \right)\end{aligned}$$

And

$$attr_i(\mathbf{x})|X, \mathbf{y} \sim \mathcal{N}(\mu'_i(x_i - \tilde{x}_i), \Sigma'_{ii}(x_i - \tilde{x}_i)^2)$$

Integrated Gradients

In this work, we use the Integrated Gradients definition of attributions:

$$attr_i(\mathbf{x}|F) = (x_i - \tilde{x}_i) \int_0^1 \frac{\partial F(\tilde{\mathbf{x}} + t(\mathbf{x} - \tilde{\mathbf{x}}))}{\partial x_i} dt$$

The **Fundamental Theorem of Line Integrals** states the following:

$$\int_{\Gamma} \nabla F \, d\mathbf{x} = F(\gamma(1)) - F(\gamma(0))$$

$$\nabla = \left[\frac{\partial}{\partial x_1} \quad \cdots \quad \frac{\partial}{\partial x_D} \right]$$

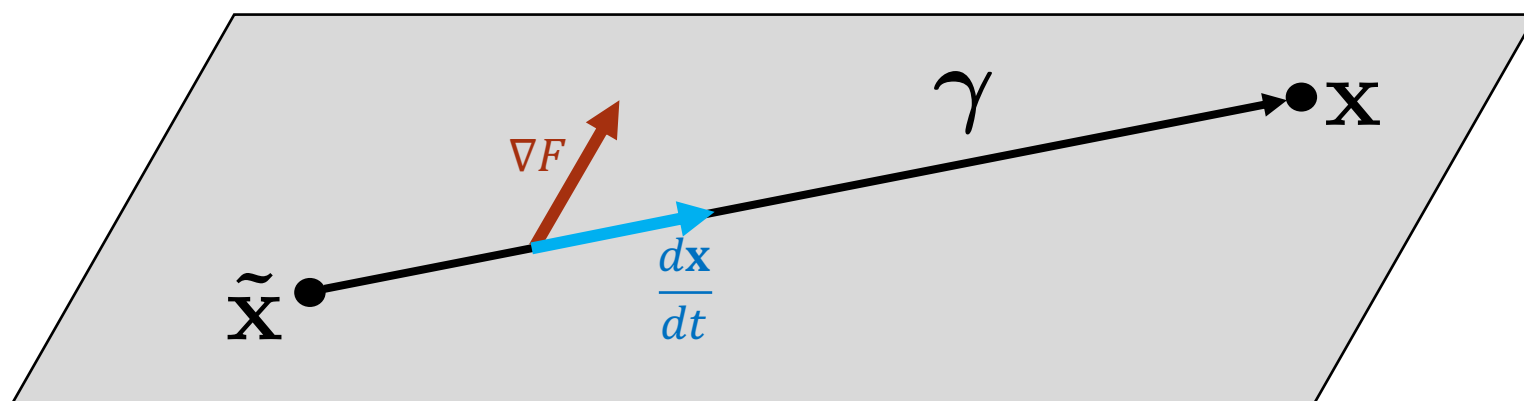
and Γ is a path in $\mathbb{R}^D$ parameterized by $\gamma(t)$, for $0 \leq t \leq 1$.

Geometric motivation for Integrated Gradients

- Let $\gamma: [0,1] \rightarrow \mathbb{R}^D$ be a curve such that $\gamma(t) = \tilde{\mathbf{x}} + t(\mathbf{x} - \tilde{\mathbf{x}})$.
- From the fundamental theorem of line integrals (Stokes' theorem):

$$F(\gamma(1)) - F(\gamma(0)) = \int_{\gamma} \nabla F \cdot d\mathbf{x}$$

$$F(\mathbf{x}) - F(\tilde{\mathbf{x}})$$



$$\int_{\Gamma} \nabla F d\mathbf{x} = \int_0^1 \sum_{i=1}^D \frac{\partial F(\gamma(t))}{\partial \gamma_i} \frac{d\gamma_i}{dt} dt$$

$$= \sum_{i=1}^D (x_i - \tilde{x}_i) \int_0^1 \frac{\partial F(\gamma(t))}{\partial \gamma_i} dt$$

$$= \sum_{i=1}^D attr_i(\mathbf{x}|F)$$

Attributions of GPs are GPs

Theorem: If F is distributed according to a GP,

$$F \sim \mathcal{GP}(m, k)$$

where $m \in C^1(\mathbb{R}^D)$ and $k \in C^2(\mathbb{R}^D \times \mathbb{R}^D)$, then

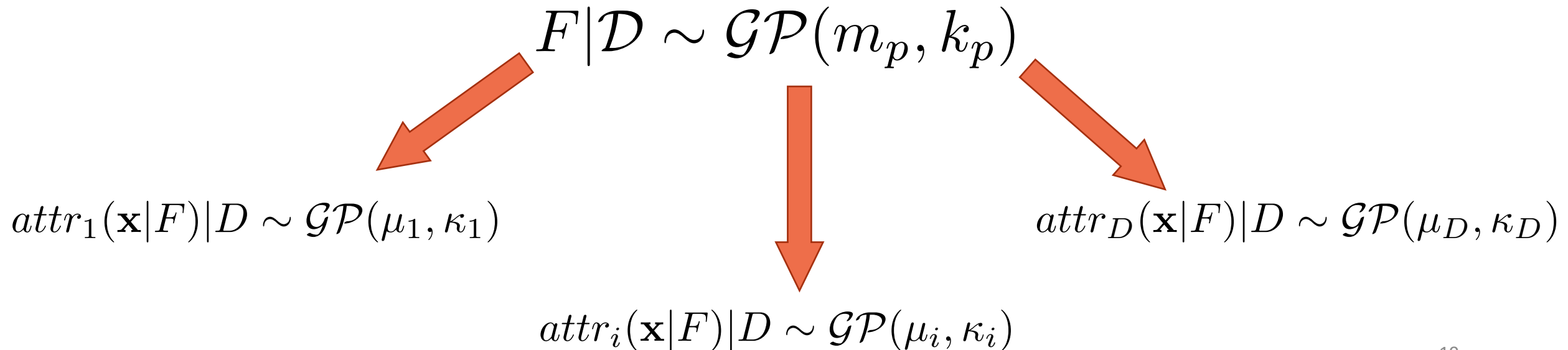
$$\text{attr}_i(\mathbf{x}|F) \sim \mathcal{GP}(\mu_i, \kappa_i)$$

where the mean and covariance functions are given by

$$\begin{aligned}\mu_i(\mathbf{x}) &= (x_i - \tilde{x}_i) \int_0^1 \frac{\partial m(\tilde{\mathbf{x}} + t(\mathbf{x} - \tilde{\mathbf{x}}))}{\partial x_i} dt, \\ \kappa_i(\mathbf{x}, \mathbf{x}') &= (x_i - \tilde{x}_i)(x'_i - \tilde{x}_i) \int_0^1 \int_0^1 \frac{\partial^2 k(\tilde{\mathbf{x}} + s(\mathbf{x} - \tilde{\mathbf{x}}), \tilde{\mathbf{x}} + t(\mathbf{x}' - \tilde{\mathbf{x}}))}{\partial x_i \partial x'_i} dt ds\end{aligned}$$

Practical result

- The theorem shows that Integrated Gradients attributions preserve Gaussianity.
- Since we obtain a GP posterior from GPR, it follows that the attributions of functions modeled by GPRs are also GPs.
- Thus, we can sensibly discuss the distribution of $attr_i(\mathbf{x}|F)$ and its properties.



Closed-Form Expressions for Attributions

- We can provide closed form expressions and code implementing the attributions of GPs with squared exponential type kernels.

$$k_{SE}(\mathbf{x}, \mathbf{x}') = \sigma_0^2 \exp \left(- \sum_{i=1}^D \frac{(x_i - x'_i)^2}{2\ell^2} \right)$$

- After some derivations, the attributions have the following mean and variances,

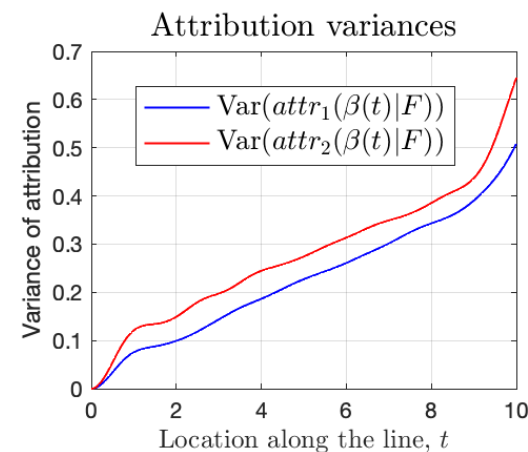
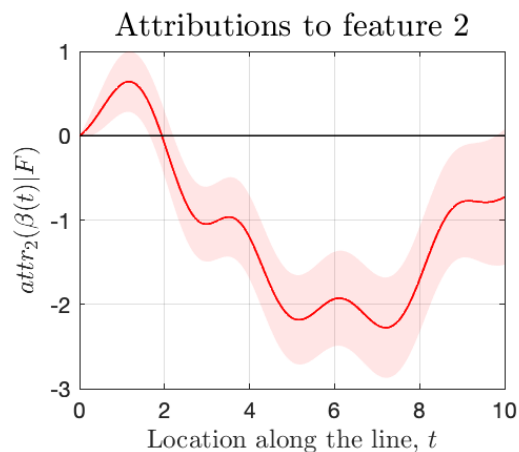
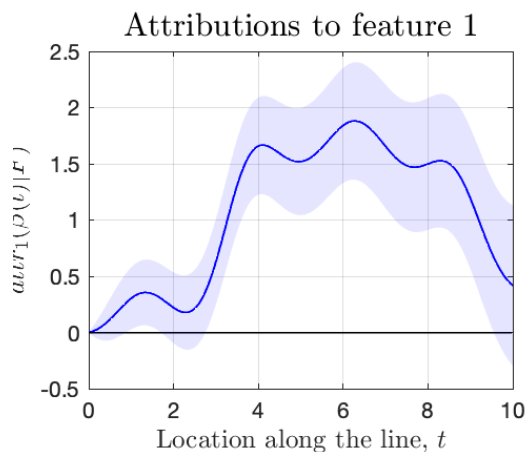
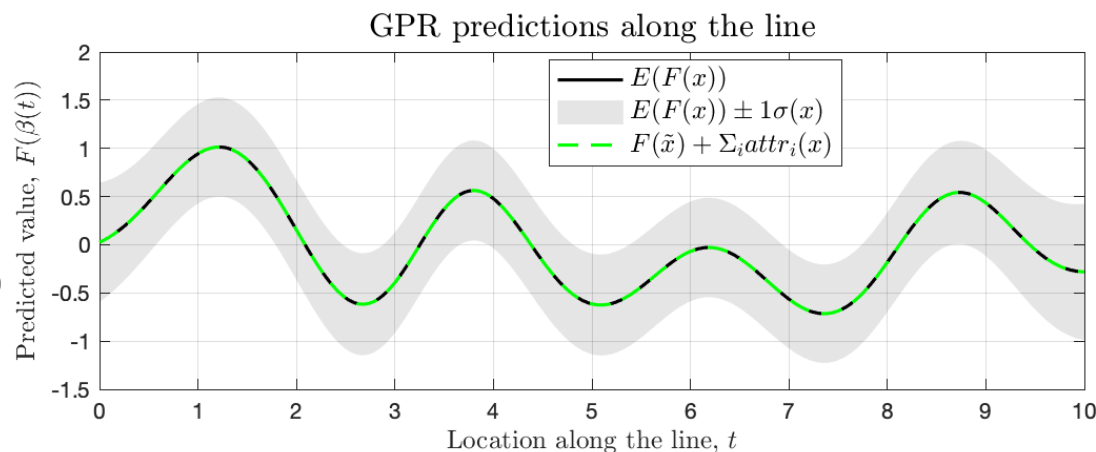
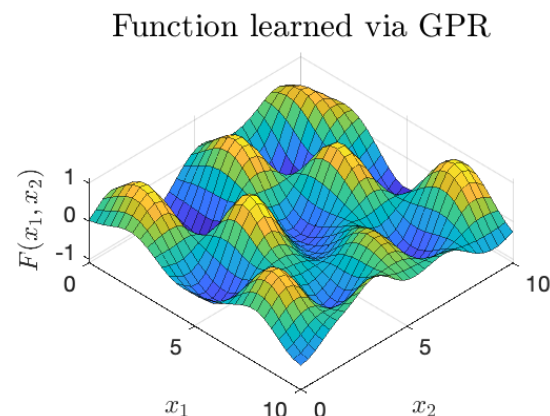
$$\mathbb{E}(\text{attr}_i(\mathbf{x})) = \sum_{n=1}^N \alpha_n A_{n,i}(\mathbf{x}),$$

$$\text{Var}(\text{attr}_i(\mathbf{x})) = B_i(\mathbf{x}) - \sum_{n=1}^N \sum_{m=1}^N \Lambda_{n,m} A_{n,i}(\mathbf{x}) A_{m,i}(\mathbf{x}),$$

Experiment with Simulated Data

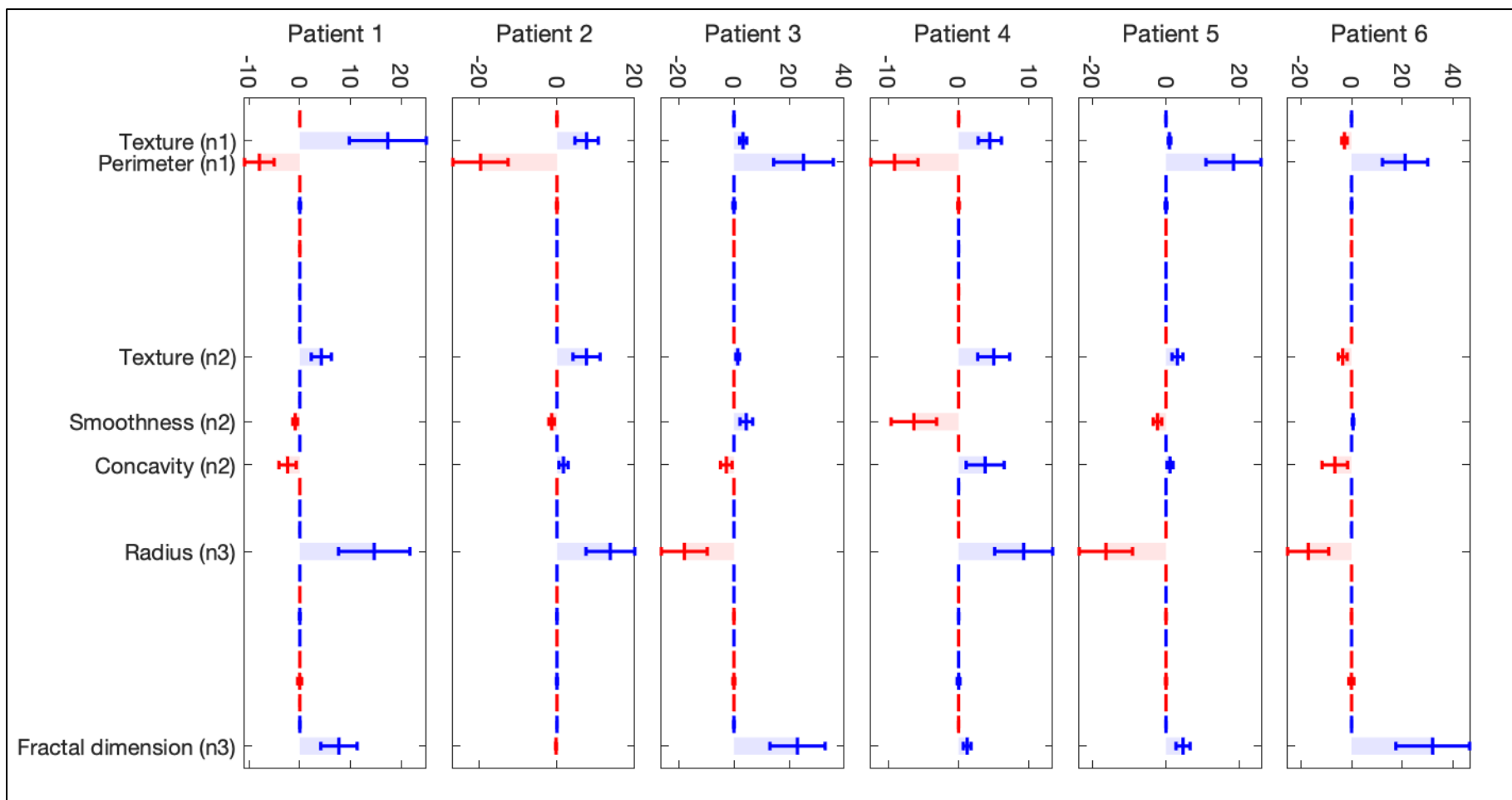
The data were generated according to $X_1 \sim \mathcal{U}(0,10), X_2 \sim \mathcal{U}(0,10)$,

$$N_Y \sim \mathcal{N}(0,1), \quad y = \sin x_1 \sin 1.5x_2 + \frac{1}{2} n_y$$



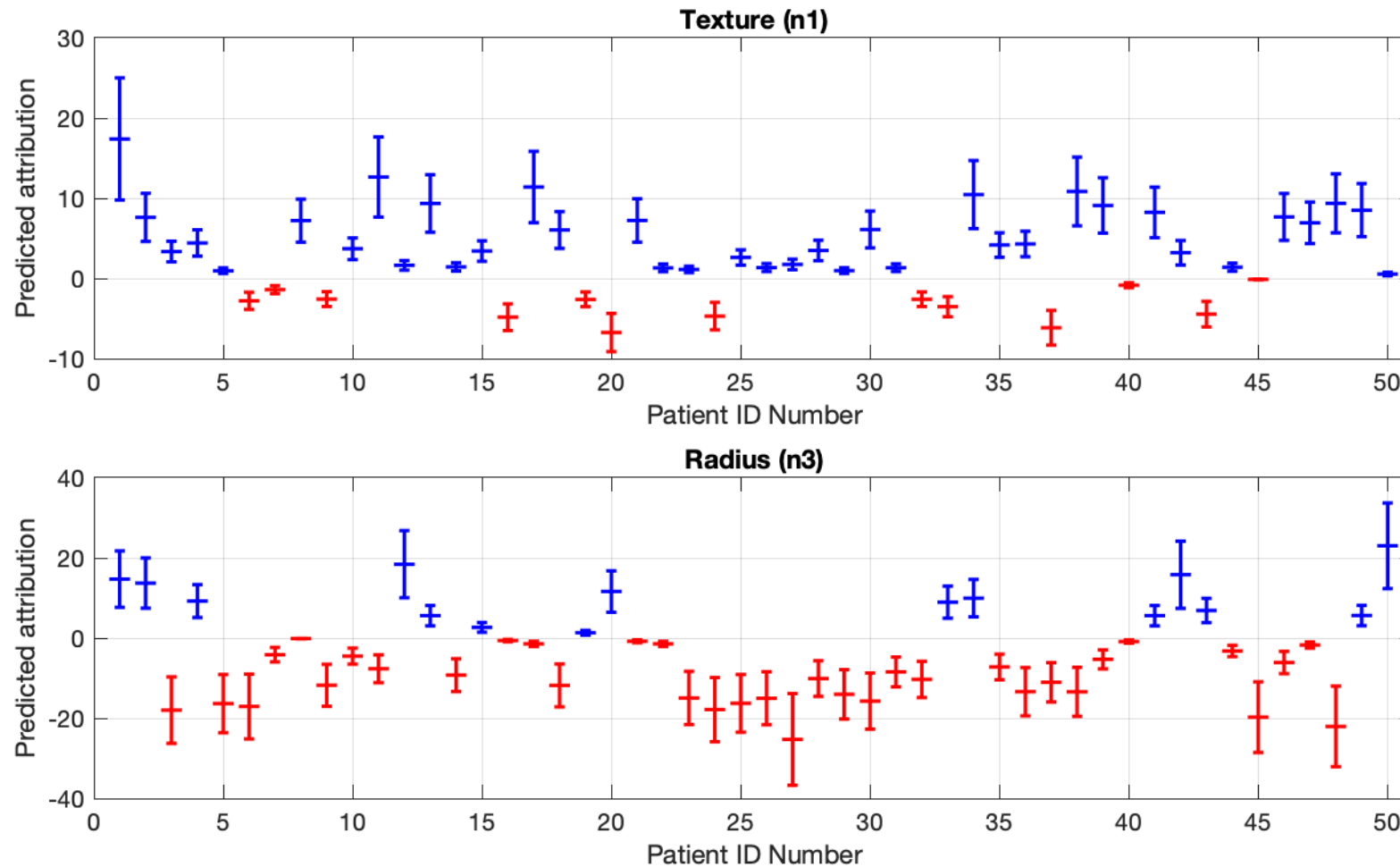
Breast Cancer Diagnosis

- GPR provides uncertainty in predictions
- More important features (of the 30 features) have more uncertainty.

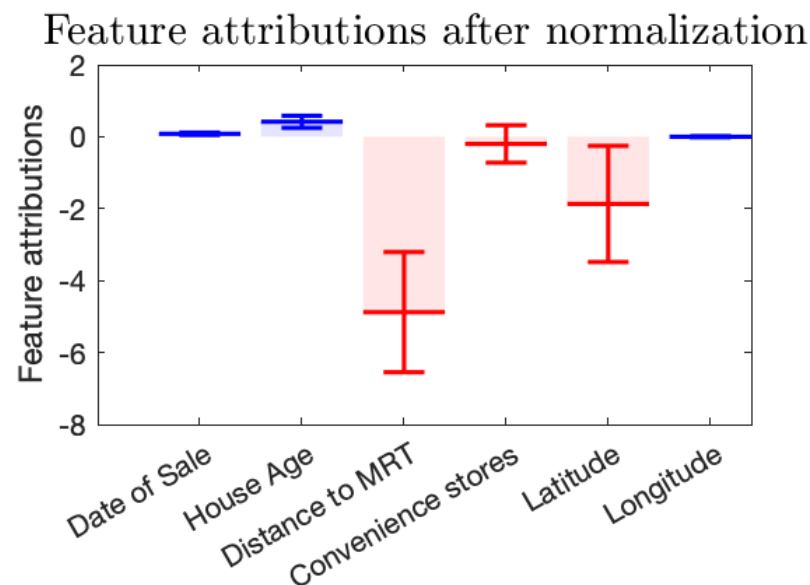
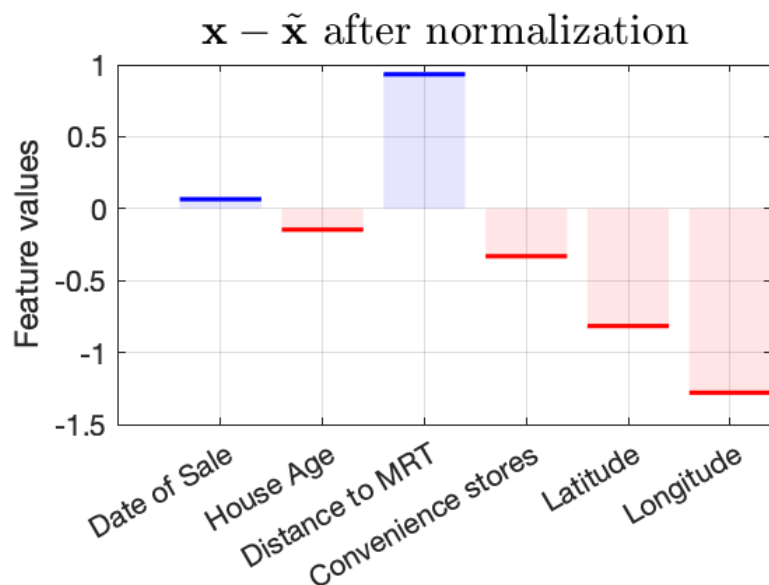
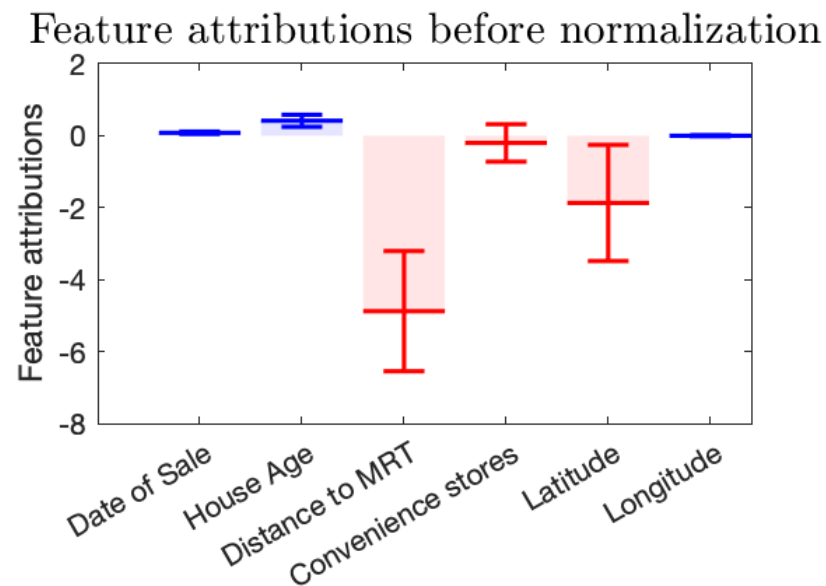
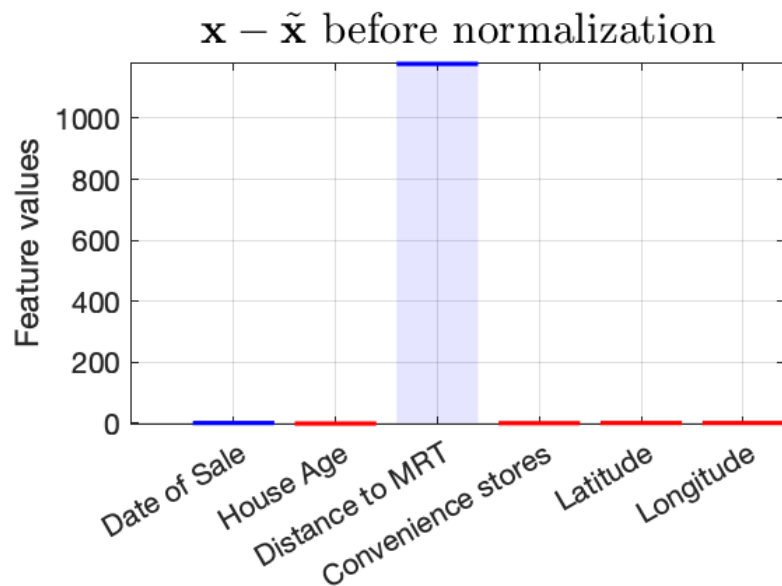


Breast Cancer Diagnosis (cont'd)

The variation of attribution of two features across patients



Predictions using the Taipei Housing Data



Random Feature GPs

- Random feature GPs use the spectral representation of the kernel function to approximate the kernel.

$$k(\mathbf{x}, \mathbf{x}') \approx \phi(\mathbf{x})^\top \phi(\mathbf{x}') \quad \phi(\mathbf{x}) = \begin{bmatrix} \sin(\mathbf{v}_1^\top \mathbf{x}) \\ \cos(\mathbf{v}_1^\top \mathbf{x}) \\ \vdots \\ \sin(\mathbf{v}_M^\top \mathbf{x}) \\ \cos(\mathbf{v}_M^\top \mathbf{x}) \end{bmatrix}$$

- The frequency vectors $\mathbf{v}_m$ are randomly sampled from the power spectral density of the kernel.
- This representation yields equations which are more tractable during the training of a GPR model.

$$\begin{aligned} \mathbb{E}(F(\mathbf{x})|\mathcal{D}) &= m(\mathbf{x}) + \mathbf{k}(\mathbf{x})^\top (\mathbf{K} + \sigma_n^2 \mathbf{I})^{-1} (\mathbf{y} - \mathbf{m}) \\ &\approx m(\mathbf{x}) + \phi(\mathbf{x})^\top \Phi (\Phi^\top \Phi + \sigma_n^2 \mathbf{I})^{-1} (\mathbf{y} - \mathbf{m}) \end{aligned}$$

Random Feature GPs

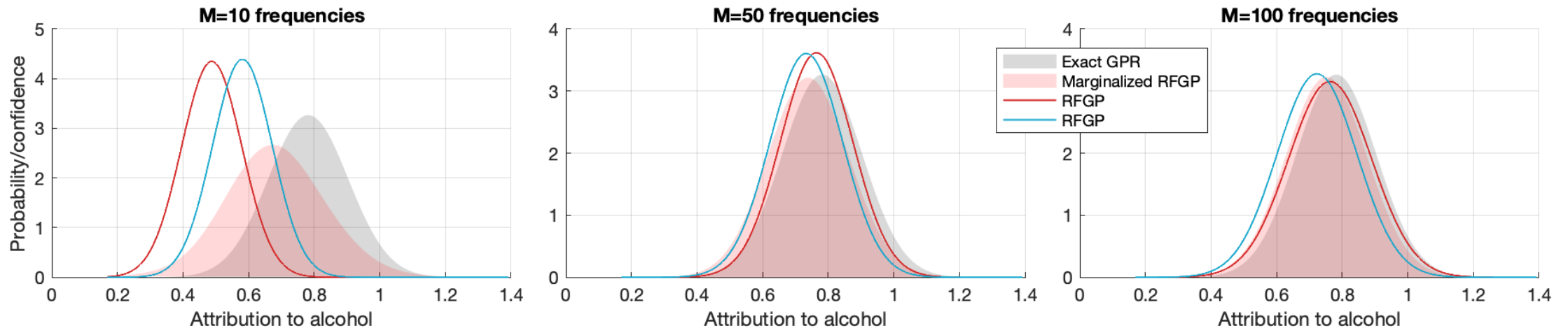
- Random feature GPs admit simple to evaluate attributions.
- Since RFGPs approximate GPR, the RFGP attributions should approximate the GPR attributions.

$$\mathbb{E}(\text{attr}_i(\mathbf{x})) = (x_i - \tilde{x}_i)\zeta(\mathbf{x})\mathbf{A}^{-1}\Phi\mathbf{y},$$

$$\text{Var}(\text{attr}_i(\mathbf{x})) = (x_i - \tilde{x}_i)^2\sigma_n^2\zeta(\mathbf{x})^\top\mathbf{A}^{-1}\zeta(\mathbf{x}),$$

$$\zeta(\mathbf{x}) = \begin{bmatrix} \frac{v_{1,i}}{\mathbf{v}_1^\top(\mathbf{x}-\tilde{\mathbf{x}})} (\cos(\mathbf{v}_1^\top\mathbf{x}) - \cos(\mathbf{v}_1^\top\tilde{\mathbf{x}})) \\ \frac{v_{1,i}}{\mathbf{v}_1^\top(\mathbf{x}-\tilde{\mathbf{x}})} (\sin(\mathbf{v}_1^\top\mathbf{x}) - \sin(\mathbf{v}_1^\top\tilde{\mathbf{x}})) \\ \vdots \\ \frac{v_{M,i}}{\mathbf{v}_M^\top(\mathbf{x}-\tilde{\mathbf{x}})} (\cos(\mathbf{v}_M^\top\mathbf{x}) - \cos(\mathbf{v}_M^\top\tilde{\mathbf{x}})) \\ \frac{v_{M,i}}{\mathbf{v}_M^\top(\mathbf{x}-\tilde{\mathbf{x}})} (\sin(\mathbf{v}_M^\top\mathbf{x}) - \sin(\mathbf{v}_M^\top\tilde{\mathbf{x}})) \end{bmatrix} \in \mathbb{R}^{2M \times 1},$$

$$\mathbf{A} = \Phi\Phi^\top + \frac{M\sigma_n^2}{k(0,0)}\mathbf{I}_{2M}$$



Conclusions

- We developed a bridge between the theory of feature attribution and GPR.
- By representing attributions as distributions, we can quantify a model's confidence in its feature attributions.
- When the selected GPR model admits analytically tractable attributions, we mitigate the need for computationally expensive or otherwise inaccurate approximations.
- While non-parametric methods like GPR, at face value, offer a non-interpretable approach to ML, our study of feature attribution may provide more trust in GPR models.
- Furthermore, these results will extend the theoretical and practical scope of GPR models to more XAI domains.