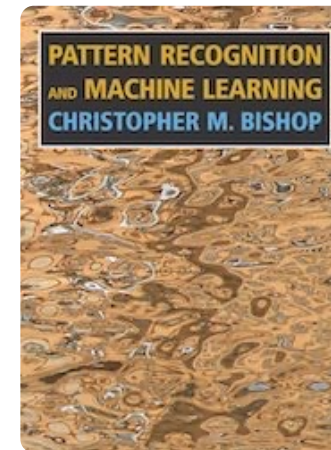
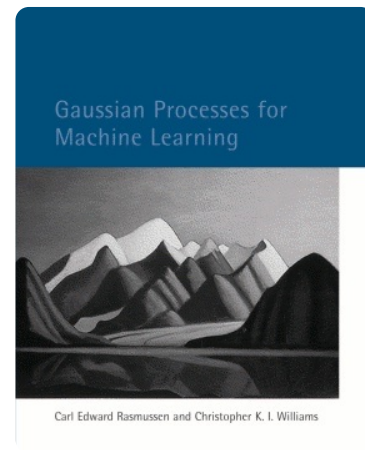


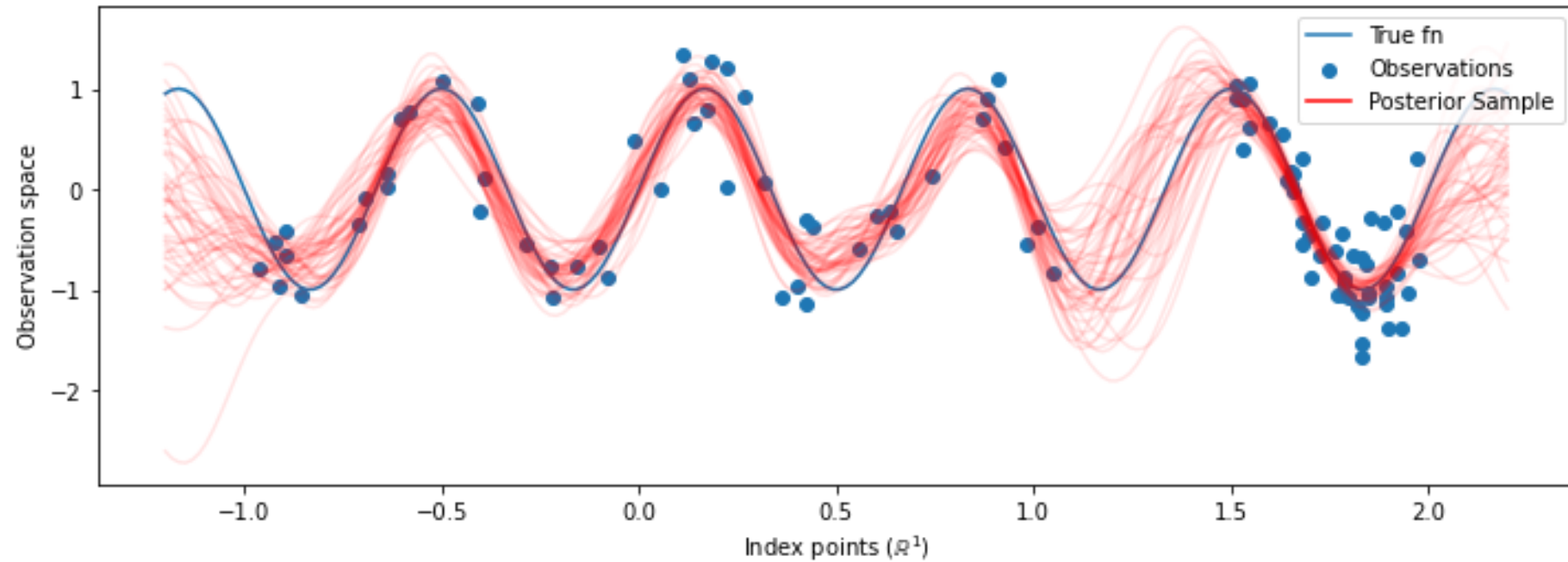
Gaussian Processes for Regression

References

- Gaussian processes for ML (MIT Press 2006) [RW]
- PRML section 6.4
- Tensorflow probability



Gaussian Processes for Regression



Quantify the uncertainty (or uncertainty quantify)

Gaussian Process to describe a distribution over functions.

Def A Gaussian process is a collection of random variables, any finite number of which have a joint Gaussian distribution [RW]

$\mathcal{X}$: any set (index)

A random process $X(t)$ is a GP indexed by $\mathcal{X}$ if for any $\{t_1, \dots, t_n\} \subset \mathcal{X}$, a random vector formed by $X(t_1), \dots, X(t_n)$ is jointly Gaussian (or $(X(t_1), \dots, X(t_n))^T$ is a multivariate Gaussian)

Specification of GP

GP can be specified by

- mean function $\mu: \mathcal{X} \rightarrow \mathbb{R}$ s.t. $\mu(t) = \mathbb{E}[X(t)]$
- covariance function $k: \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$ s.t.

$$K(t, s) = \text{cov}[X(t), X(s)] \quad \text{a.k.a kernel function}$$

The kernel function is required only to be symmetric and positive definite.

$$X(t) \sim \text{GP}(\mu(t), K(t, s)) \quad t, s \in \mathcal{X}$$

Example D -dim multivariate Gaussian distribution is a

GP indexed by $\{1, 2, \dots, D\}$

See PRML section 2.3.2

Example $X(t) := tx$ where $x \sim \mathcal{N}(x | 0, 1)$ and $t \in \mathbb{R}$

Then $X(t)$ is a GP indexed by $\mathbb{R}$

$$\mu(t) = \mathbb{E}[X(t)] = t \mathbb{E}[x] = 0$$

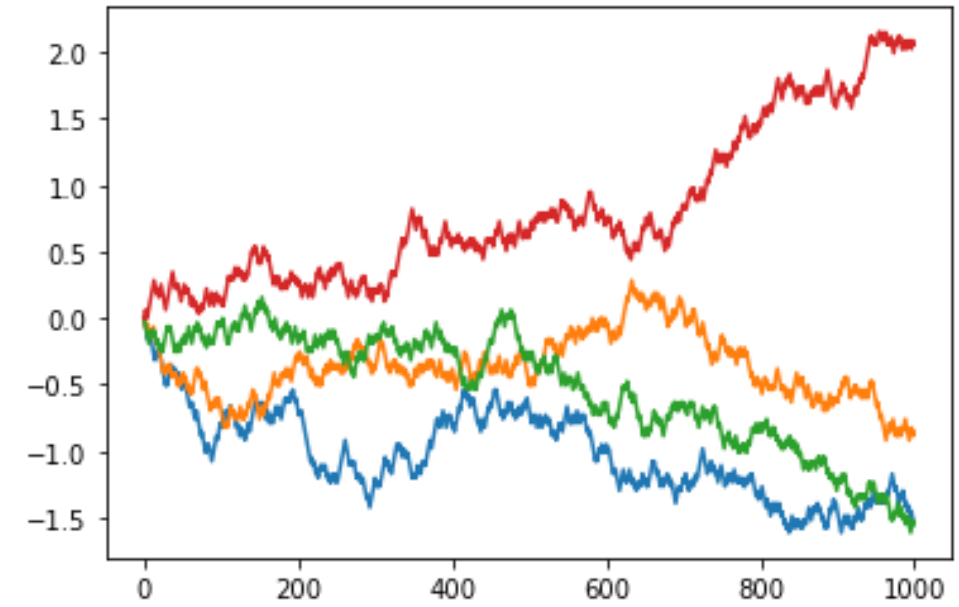
$t_1, \dots, t_n, \quad x(t_1), \dots, x(t_n)$

$$K(t, s) = \text{cov}[X(t), X(s)] = \mathbb{E}[tx \cdot sx] = ts \mathbb{E}[x^2] = ts$$

Def A stochastic process $\{X(t), t \geq 0\}$ is said to be
(random)

a Brownian motion (Wiener process) if

- $X(0) = 0$
- $\{X(t), t \geq 0\}$ has stationary and independent increments
- For $t > 0$, $X(t)$ is normally distributed with mean 0 and variance t .
- With probability 1, $t \mapsto X(t)$ is continuous.



For $t_1 < t_2 < \dots < t_n$, joint density $X(t_1), \dots, X(t_n)$

$X(t_1), X(t_2) - X(t_1), \dots, X(t_n) - X(t_{n-1})$ are independent and

$$X(t_k) - X(t_{k-1}) \sim N(0, t_k - t_{k-1})$$

Eg. Brownian motion $\{B(t), t \geq 0\}$ is a GP indexed by $\mathbb{R}^+$

Each of $B(t_1) \dots B(t_n)$ can be expressed as a linear

combination of independent normal r.v $B(t_1), B(t_2) - B(t_1) \dots, B(t_n) - B(t_{n-1})$

Gaussian Processes for Regression (GPR)

Notations

$\mathcal{X}$: index set of GP

$\mathbf{z}$: collection of random variables indexed by $\mathcal{X}$

$\mathbf{z}(x) \sim \text{GP}(\mu(x), k(x, x'))$ $x, x' \in \mathcal{X}$

For any finite set $x_1, x_2, \dots, x_n \in \mathcal{X}$,

$$(z(x_1), \dots, z(x_n))^T = \mathcal{N}(\mu_n, K)$$

where $\mu_n = (\mu(x_1), \mu(x_2), \dots, \mu(x_n))^T$ $K_{ij} = K(x_i, x_j)$
 $= \text{cov}[z(x_i), z(x_j)]$

index set of GP $\approx$ domain of function

GPR : nonparametric Bayesian regression method using
the properties of GP

Gaussian Processes for Regression

Goal : estimate an unknown function f with noisy observations $\{t_1, \dots, t_N\}$ of f at points $\{x_1, \dots, x_N\}$ with error bars (u.q.)

Idea : GP as defining a distribution over functions and inference taking place directly in the space of functions.

We dispense with the parametric model and instead define a prior distribution over functions directly.

Linear regression revisited

Training data set $D = (\mathbf{x}_n, t_n)_{n=1 \dots N}$

$\mathbf{x}$: D -dim input vector, t : target value

$y(\mathbf{x})$: prediction value at $\mathbf{x}$

predictive distribution

Goal: find $p(t^* | \mathbf{x}^*, \mathcal{X}, \mathcal{T})$ at some new input $\mathbf{x}^*$

Fix basis functions $\phi_1, \dots, \phi_{M-1}$ or feature map $\Phi(\mathbf{x})$

Linear model

$$y(\mathbf{x}, \mathbf{w}) := \mathbf{w}^T \Phi(\mathbf{x}) \quad \text{weight vector } \mathbf{w} \in \mathbb{R}^M$$

Model assumption

$$t \sim \mathcal{N}(t | y(x, w), \beta^{-1})$$

precision $\beta > 0$

prior over w

$$P(w | \alpha) = \mathcal{N}(w | \theta, \alpha^{-1} I)$$

posterior over w

$$P(w | \mathbb{t}) = \mathcal{N}(w | m_N, S_N)$$

$$m_N = \beta S_N \Phi^T \mathbb{t}$$

$$S_N^{-1} = \alpha I + \beta \Phi^T \Phi$$

Example

Let $w_l \sim N(w_l | \theta, \alpha^{-1} I)$ be as a prior.

For $n=1, 2, \dots, N$, let $y_n := y(x_n, w_l) = w_l^T \Phi(x_n)$. Then y_n can be seen a 1-dim random variable (1-dim Gaussian)

Let $\mathbf{y} := (y_1 \dots y_N)^T$. Then

$$\mathbf{y} = \overset{N \times M}{\Phi} \underset{M \times 1}{w_l}$$

is the N -dim random vector where $\Phi_{nk} = \phi_k(x_n)$.

Φ is called design matrix so that

$$\Phi = \begin{pmatrix} \phi_0(x_1) & \phi_1(x_1) & \cdots & \phi_{M-1}(x_1) \\ \phi_0(x_2) & \phi_1(x_2) & & \phi_{M-1}(x_2) \\ \vdots & \vdots & & \\ \phi_0(x_N) & & & \phi_{M-1}(x_N) \end{pmatrix} = \begin{pmatrix} \Phi(x_1)^T \\ \Phi(x_2)^T \\ \vdots \\ \Phi(x_N)^T \end{pmatrix}$$

Since $\mathbf{y}$ is a linear combination of w_i (Gaussian dist.),
 $\mathbf{y}$ is N -dim Gaussian. ($\mathbf{y}$ is determined by the first
and the second moments)

$$\mathbb{E}[\mathbf{y}] = \mathbb{E}[\Phi \mathbf{w}] = \Phi \mathbb{E}[\mathbf{w}] = \mathbf{0}$$

$$\text{Cov}[\mathbf{y}] = \mathbb{E}[\mathbf{y} \mathbf{y}^T] = \Phi \underbrace{\mathbb{E}[\mathbf{w} \mathbf{w}^T]}_{\sigma^2 \mathbf{I}} \Phi^T =: \mathbf{K}$$

where $\mathbf{K}$ is the gram matrix

$$K_{nm} = K(\mathbf{x}_n, \mathbf{x}_m) = \frac{1}{\alpha} \Phi(\mathbf{x}_n)^T \Phi(\mathbf{x}_m)$$

and $K(\mathbf{x}, \mathbf{x}')$ is the kernel function.

In general, GP is defined as a probability distribution over functions $y(x)$ s.t the set of values of $y(x)$ evaluated at an arbitrary set $x_1, \dots, x_N$ jointly have a Gaussian Distribution.

A key point about Gaussian processes is that the joint distribution over N variables $y_1 \dots y_N$ is specified completely by the second-order statistics.

①

Since we have not any knowledge about the mean of $y(x)$, we take it to be zero. This is equivalent

to choosing the mean vector of prior $p(w|\alpha)$ to be $\mathbf{0}$

②

The specification of GP is then completed by giving the covariance of $y(x)$ evaluated at any two values of x with kernel $\mathbb{E}[y(x_n), y(x_m)] = k(x_n, x_m)$.

Sketch of GPR

Let $f(x)$ ($y(x)$) be a zero mean GP indexed by its domain space.

I.e.

$$\begin{pmatrix} y(x_1) \\ \vdots \\ y(x_N) \end{pmatrix} \sim \mathcal{N} \left(\mathbf{0}, \begin{bmatrix} k(x_1, x_1) & \cdots & k(x_1, x_N) \\ \vdots & \ddots & \vdots \\ k(x_N, x_1) & \cdots & k(x_N, x_N) \end{bmatrix} \right)$$

and

$$\begin{pmatrix} y(x_1) \\ \vdots \\ y(x_N) \\ y(x^*) \end{pmatrix} \sim \mathcal{N} \left(\mathbf{0}, \begin{bmatrix} k(x_1, x_1) & \cdots & k(x_1, x_N) & k(x_1, x^*) \\ \vdots & \ddots & \vdots & \vdots \\ k(x_N, x_1) & \cdots & k(x_N, x_N) & \vdots \\ k(x^*, x_1) & \cdots & k(x^*, x_N) & k(x^*, x^*) \end{bmatrix} \right)$$

$\mathbb{K}^*$

Assume $t = y + \varepsilon_n$ with $\varepsilon_{\text{noise}} \sim \mathcal{N}(0, \beta^{-1})$

$\Rightarrow$

$$\begin{pmatrix} t_1 \\ \vdots \\ t_N \end{pmatrix} \sim \mathcal{N}(\mathbf{0}, K + \beta^{-1} I_N)$$

and

$$\begin{pmatrix} t_1 \\ \vdots \\ t_N \\ t^* \end{pmatrix} \sim \mathcal{N}(\mathbf{0}, K^* + \beta^{-1} I_{N+1})$$

Since $p(\mathbb{t}, t^*)$ is Gaussian so is $p(t^* | \mathbb{t})$

$$p(t^* | \mathbb{t}) = \mathcal{N}(t^* | m(x^*), \sigma^2(x^*))$$

//

GPR in detail

We shall consider the noise on the observed target value

Assumption: Gaussian noisy model.

$$t_n = y_n + \epsilon_n \quad n = 1, 2, \dots, N$$

where ϵ_n followed zero mean Gaussian with precision β .

So for an N -dim vector $\mathbf{y} := (y_1 \dots y_N)^T$, the probability distribution over $\mathbf{t} := (t_1, \dots, t_N)^T$ can be expressed as

$$P(\mathbf{t} | \mathbf{y}) = \mathcal{N}(\mathbf{t} | \mathbf{y}, \beta^{-1} \mathbf{I}_N)$$

where I_N is the $N \times N$ identity matrix.

From the definition of GP, (prior dist. over $\mathbf{y}$)

$$p(\mathbf{y}) = \mathcal{N}(\mathbf{y} | \mathbf{0}, \mathbf{K}) \quad N\text{-dim Gaussian}$$

The kernel function determining the gram matrix $\mathbf{K}$ is typically chosen to express the property that

$$\mathbf{x}_n \sim \mathbf{x}_m \Rightarrow y(\mathbf{x}_n) \text{ and } y(\mathbf{x}_m) \text{ are strongly correlated}$$

One widely used kernel function for GPB is given by

$$K(\mathbf{x}_n, \mathbf{x}_m) := \sigma_0 \exp \left\{ -\frac{\sigma_1}{2} \|\mathbf{x}_n - \mathbf{x}_m\|^2 \right\} + \sigma_2 + \sigma_3 \mathbf{x}_n^T \mathbf{x}_m$$

Now we consider the distribution over $\mathbf{z}$. By PRML section 2.3.3 (linear Gaussian model), we obtain

$$p(\mathbf{z}) = \int p(\mathbf{z} | \mathbf{y}) p(\mathbf{y}) d\mathbf{y} = \mathcal{N}(\mathbf{z} | \boldsymbol{\mu}, \mathbb{C}) \quad (6.61)$$

where $\mathbb{C} = \mathbf{K} + \beta^{-1} \mathbf{I}_N$ with $C(\mathbf{x}_n, \mathbf{x}_m) = K(\mathbf{x}_n, \mathbf{x}_m) + \beta^{-1} \delta_{nm}$

Using Gaussian process viewpoint, we shall build a model of the joint probability distribution over sets of data points.

Goal : Given training data points, make prediction of target value for new input

x^* : some new input, t^* : its target value

(X, t) set of training data

Find the predictive distribution $p(t^* | t) = p(t^* | x^*, X, t)$

Let $\mathbf{t}^* := (t_1, \dots, t_N, t^*)^T$ ($N+1$ - dim random vector)

To find $p(t^* | \mathbf{t})$, we first consider $p(\mathbf{t}^*)$. From (6.61),

$$p(\mathbf{t}^*) = \mathcal{N}(\mathbf{t}^* | \mathbf{0}, \mathbb{C}^*) \quad \begin{array}{l} N+1 \text{ - dim} \\ \text{Gaussian} \end{array}$$

The covariance matrix $\mathbb{C}^*$ can be decomposed as

$$\mathbb{C}^* = \begin{pmatrix} \overset{N \times N}{\mathbb{C}} & \mathbf{k} \\ \mathbf{k}^T & c \end{pmatrix} \quad (N+1) \times (N+1)$$

where $\mathbf{k}$ is the N -dim vector whose n^{th} element is $k(x_n, \mathbf{x}^*)$ and $c = k(\mathbf{x}^*, \mathbf{x}^*) + \beta^{-1}$.

Since the conditional Gaussian distribution is again a Gaussian, $p(t^* | \mathbf{z})$ is Gaussian.

$$\mathbf{z}^* = \begin{pmatrix} \mathbf{z} \\ t^* \end{pmatrix}$$

From (2.81) and (2.82),

$N+1$ dim

$$p(t^* | \mathbf{z}) = \mathcal{N}(t^* | m(\mathbf{z}^*), \sigma^2(\mathbf{z}^*))$$

where

$$m(\mathbf{z}^*) = \mathbf{K}^T \mathbf{C}^{-1} \mathbf{z}$$

$$\sigma^2(\mathbf{z}^*) = c - \mathbf{K}^T \mathbf{C}^{-1} \mathbf{K}$$

) depending on $\mathbf{z}^*$

The expectation of predictive distribution can be written as

$$m(x^*) = \sum_{n=1}^N a_n k(x_n, x^*)$$

where a_n is the n^{th} component of $C^{-1} \mathbb{t}$

Remark :

- Computational costs of GPR: the inversion of $N \times N$ matrix requires $O(N^3)$ once. For each new point, GPR has the cost $O(N^2)$ from the matrix multiplication

GPR in the real world (learning for hyperparameters)

Hyperparameter $\theta := \{ \theta_0, \theta_1, \theta_2, \theta_3 \}$ and β

Log likelihood function

$$\ln P(\mathbf{t} | \theta) = -\frac{1}{2} \ln |\mathbb{C}| - \frac{1}{2} \mathbf{t}^T \mathbb{C}^{-1} \mathbf{t} - \frac{N}{2} \ln(2\pi)$$

$$\frac{\partial}{\partial \theta_i} \ln P(\mathbf{t} | \theta) = -\frac{1}{2} \left(\mathbb{C}^{-1} \frac{\partial \mathbb{C}}{\partial \theta_i} \right) + \frac{1}{2} \mathbf{t}^T \mathbb{C}^{-1} \frac{\partial \mathbb{C}}{\partial \theta_i} \mathbb{C}^{-1} \mathbf{t}$$

Generally, $\ln P(\mathbf{t} | \theta)$ is not convex.

Tensorflow Probability

Exponentiated Quadratic

$$K(x_n, x_m) := \sigma_0^2 \exp \left(\frac{\|x_n - x_m\|^2}{-2 \sigma_1^2} \right)$$

σ_0 : amplitude

σ_1 : length scale

$1/\beta$: observation noise variance.

$$K(x_n, x_m) := \sigma_0 \exp \left\{ -\frac{\sigma_1}{2} \|x_n - x_m\|^2 \right\} + \sigma_2 + \sigma_3 x_n^T x_m$$