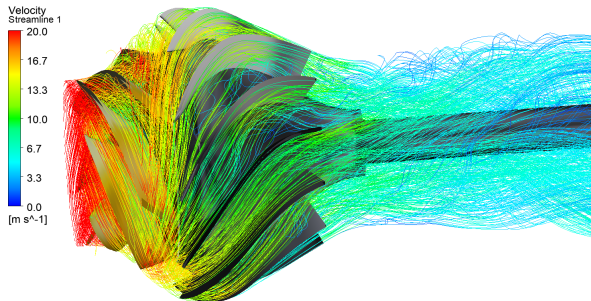


Data Driven Approaches in (Multi-objective) Bayesian Optimisation - An Introductory Course

Tinkle Chugh

Department of Computer Science
University of Exeter, UK

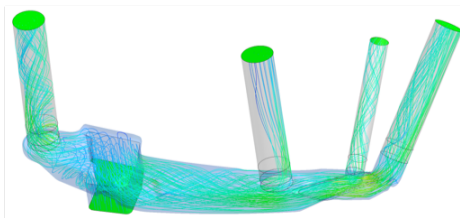
Expensive Optimisation



Pump Optimisation

- 1 One evaluation: 24 to 48 hours
- 2 Utilises CFD simulations

Expensive Optimisation



Air Intake Ventilation System

- 1 One evaluation: 30 minutes
- 2 Utilises CFD simulations

Bayesian Optimisation

Problem

$$\min_{\mathbf{x} \in \mathcal{X}} f(\mathbf{x}) \quad \mathbf{x} \in \mathcal{X} \subset \mathbb{R}^d$$

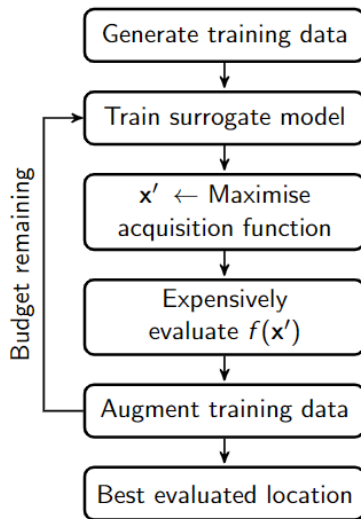
Bayesian optimisation

- ▶ Build a probabilistic surrogate model of $f(\mathbf{x})$
- ▶ Maximise an *acquisition* function $\alpha(\mathbf{x})$

$$\mathbf{x}' \leftarrow \operatorname{argmax}_{\mathbf{x} \in \mathcal{X}} \alpha(\mathbf{x})$$

Balances exploitation with exploration

- ▶ Expensively evaluate $f(\mathbf{x}')$



Two important elements

- ① Bayesian model: $p(f|D)$
 - ② Acquisition function : $\alpha(x, p(f|D))$
- Focus of this lecture: Bayesian Model

Gaussian Processes

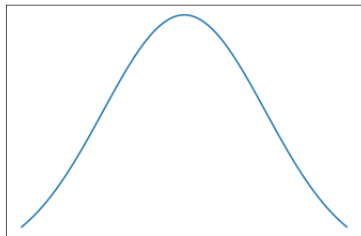
- Time series: Wiener, Kolmogorov 1940's
- Geostatistics: kriging 1970's
- Spatial statistics in general: see Cressie [1993] for overview
- General regression: O'Hagan [1978]
- Computer experiments (noise free): Sacks et al. [1989] and a survey on Bayesian optimization: Shahriari et al. [2016]
- Machine learning: Williams and Rasmussen [1996], Neal [1996]

Reading material: Carl Edward Rasmussen and Christopher K. I. Williams. Gaussian Processes for Machine Learning, MIT Press, 2006

<http://www.gaussianprocess.org/gpml/chapters/>

- Given a data set $D = \{(\mathbf{x}_n, y_n)\}_{n=1}^N$, we want to learn a function f , which not only provides the mean prediction (or point prediction) but also the uncertainty around them
- The prediction can be for regression or classification

Bayesian modelling



- 1 Univariate normal (or Gaussian) distribution has two parameters: mean (μ) and standard deviation (σ)

- 1 Univariate normal (or Gaussian) distribution has two parameters: mean (μ) and standard deviation (σ)
- 2 Let us denote the parameters with $\theta = (\mu, \sigma)$
- 3 Given a data $\mathbf{y} = (y_1, \dots, y_n)$, the distribution is denoted as:

$$p(y_i|\theta) \sim \mathcal{N}(\mu, \sigma^2) \text{ Or}$$

$$p(y_i|\theta) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-0.5 \frac{(y_i - \mu)^2}{\sigma^2}\right)$$

A statistical model can specify a probability distribution for data $\mathbf{y} = (y_1, y_2, \dots, y_n)$ as $p(\mathbf{y}|\theta)$

Likelihood

- $p(\mathbf{y}|\boldsymbol{\theta})$ with varying values of model parameters $\boldsymbol{\theta}$

Likelihood

- $p(\mathbf{y}|\boldsymbol{\theta})$ with varying values of model parameters $\boldsymbol{\theta}$
- In other words, how likely is it to get $\mathbf{y}$ for different values of $\boldsymbol{\theta}$

Likelihood

- $p(\mathbf{y}|\boldsymbol{\theta})$ with varying values of model parameters $\boldsymbol{\theta}$
- In other words, how likely is it to get $\mathbf{y}$ for different values of $\boldsymbol{\theta}$
- For n independent and identically distributed (IID) samples, $(y_1, y_2, \dots, y_n)$, the likelihood is:

$$L(\boldsymbol{\theta}) = \prod_{i=1}^n p(y_i|\boldsymbol{\theta})$$

$$p(\mathbf{y}|\boldsymbol{\theta}) = \prod_{i=1}^n p(y_i|\boldsymbol{\theta})$$

Likelihood

- $p(\mathbf{y}|\boldsymbol{\theta})$ with varying values of model parameters $\boldsymbol{\theta}$
- In other words, how likely is it to get $\mathbf{y}$ for different values of $\boldsymbol{\theta}$
- For n independent and identically distributed (IID) samples, $(y_1, y_2, \dots, y_n)$, the likelihood is:

$$L(\boldsymbol{\theta}) = \prod_{i=1}^n p(y_i|\boldsymbol{\theta})$$
$$p(\mathbf{y}|\boldsymbol{\theta}) = \prod_{i=1}^n p(y_i|\boldsymbol{\theta})$$

- Likelihood depends on the model - two different models for the same data $\mathbf{y}$ will typically have two different likelihoods

Maximising likelihood

- Find the model parameters θ by looking at the data $\mathbf{y}$ Or
- Maximise the likelihood to find the model parameters i.e.

$$\hat{\theta} = \underset{\theta}{\operatorname{argmax}} L(\theta) \text{ Or}$$

$$\hat{\theta} = \underset{\theta}{\operatorname{argmax}} \prod_{i=1}^n p(y_i | \theta)$$

Maximising Likelihood

- argmax of $L(\theta)$ is the same as the argmax of $\log L(\theta)$ - makes the maths simpler:

$$\log L(\theta) = \log \prod_{i=1}^n p(y_i|\theta) = \sum_{i=1}^n \log p(y_i|\theta)$$

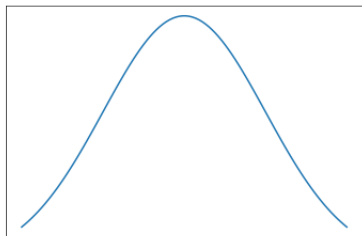
$$\hat{\theta} = \operatorname{argmax}_{\theta} \log L(\theta) = \operatorname{argmax}_{\theta} \sum_{i=1}^n \log p(y_i|\theta)$$

Maximising Likelihood

Let us assume the data is normally distributed

$$p(y_i|\theta) \sim \mathcal{N}(\mu, \sigma^2) \text{ for } i = 1, \dots, n$$

Maximise the likelihood to get the parameters



Maximising likelihood

$$p(y_i | \theta) \sim \mathcal{N}(\mu, \sigma^2)$$
$$= \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{1}{2} \frac{(y_i - \mu)^2}{\sigma^2}\right)$$


$$L(\theta) = p(\mathbf{y} | \theta) = \prod_{i=1}^n p(y_i | \theta) \quad \text{for } i=1, \dots, n$$

Maximising likelihood

$$L(\theta) = \prod_{i=1}^n \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(y_i - \mu)^2}{2\sigma^2}\right)$$

$$\hat{\theta} = \underset{\theta}{\operatorname{argmax}} L(\theta)$$

$$\hat{\theta} = (\hat{\mu}, \hat{\sigma})$$

$$\log L(\theta) = -\frac{n}{2} \log(2\pi\sigma^2) - \sum_{i=1}^n \frac{(y_i - \mu)^2}{2\sigma^2}$$

$$\frac{\partial \log L(\theta)}{\partial \mu} = 0$$

$$\frac{\partial \log L(\theta)}{\partial \sigma} = 0$$

Maximising likelihood

$$\frac{\partial}{\partial \mu} \log L(\theta) = \sum_{i=1}^n \frac{y_i - \mu}{\sigma^2} = 0$$

$$\sum_{i=1}^n y_i - \sum_{i=1}^n \mu = 0$$

$$\mu = \frac{1}{n} \sum_{i=1}^n y_i$$

$$\frac{\partial}{\partial \sigma} \log L(\theta) = -\frac{n}{2} \frac{\partial}{\partial \sigma} \log \sigma^2 - \frac{1}{2} \frac{\partial}{\partial \sigma} \left(\frac{1}{\sigma^2} \sum (y_i - \mu)^2 \right)$$

Maximising likelihood

$$= -\frac{n}{2} \frac{1}{\sigma^2} - \frac{1}{2} \left(-\frac{2}{\sigma^3} \right) \sum_{i=1}^n (y_i - \mu)^2$$

$$= -\frac{n}{\sigma} + \frac{1}{\sigma^3} \sum_{i=1}^n (y_i - \mu)^2 = 0$$

$$\frac{\sum_{i=1}^n (y_i - \mu)^2}{\sigma^3} = \frac{n}{\sigma}$$

$$\sigma = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \mu)^2} \leftarrow$$

- The parameters θ can also be uncertain

- The parameters θ can also be uncertain
 - The probability of raining tomorrow is 0.9
 - The probability of getting heads is 0.5
 - There is a one-in-a-thousand chance that Austria will win the next World Cup

- The parameters θ can also be uncertain
 - The probability of raining tomorrow is 0.9
 - The probability of getting heads is 0.5
 - There is a one-in-a-thousand chance that Austria will win the next World Cup
- Uncertainty quantification on θ using a prior distribution: $p(\theta)$

Let us assume the data $\mathbf{y} = (y_1, \dots, y_n)$ is normally distributed (let us assume σ is given and $\boldsymbol{\theta} = \mu$) i.e.

$$p(y_i|\mu) \sim \mathcal{N}(\mu, \sigma^2) \text{ for } i = 1, \dots, n$$

Let us assume the data $\mathbf{y} = (y_1, \dots, y_n)$ is normally distributed (let us assume σ is given and $\boldsymbol{\theta} = \mu$) i.e.

$$p(y_i|\mu) \sim \mathcal{N}(\mu, \sigma^2) \text{ for } i = 1, \dots, n$$

We also assume that the parameter μ also follows a normal distribution - Prior (or belief before looking at the data) i.e.

$$p(\mu) \sim \mathcal{N}(\mu_0, \sigma_0^2)$$

Let us assume the data $\mathbf{y} = (y_1, \dots, y_n)$ is normally distributed (let us assume σ is given and $\boldsymbol{\theta} = \mu$) i.e.

$$p(y_i|\mu) \sim \mathcal{N}(\mu, \sigma^2) \text{ for } i = 1, \dots, n$$

We also assume that the parameter μ also follows a normal distribution - Prior (or belief before looking at the data) i.e.

$$p(\mu) \sim \mathcal{N}(\mu_0, \sigma_0^2)$$

How to find the $p(\mu|\mathbf{y})$?

- Likelihood: $p(\mathbf{y}|\boldsymbol{\theta})$: probability that any $y = (y_1, \dots, y_n)$ will occur for given $\boldsymbol{\theta}$

Bayesian Inference

- Likelihood: $p(\mathbf{y}|\boldsymbol{\theta})$: probability that any $y = (y_1, \dots, y_n)$ will occur for given $\boldsymbol{\theta}$
- Prior: $p(\boldsymbol{\theta})$: randomness or uncertainty on the parameter $\boldsymbol{\theta}$ before observing any data

Bayesian Inference

- Likelihood: $p(\mathbf{y}|\boldsymbol{\theta})$: probability that any $y = (y_1, \dots, y_n)$ will occur for given $\boldsymbol{\theta}$
- Prior: $p(\boldsymbol{\theta})$: randomness or uncertainty on the parameter $\boldsymbol{\theta}$ before observing any data
- Posterior: $p(\boldsymbol{\theta}|\mathbf{y})$: complete probability distribution which can quantify the uncertainty in the parameters $\boldsymbol{\theta}$

Bayesian Inference

Bayes' rule:

$$p(\theta|\mathbf{y}) = \frac{p(\mathbf{y}|\theta)p(\theta)}{p(\mathbf{y})}$$

Bayes' rule:

$$p(\theta|\mathbf{y}) = \frac{p(\mathbf{y}|\theta)p(\theta)}{p(\mathbf{y})}$$

The marginal distribution of $\mathbf{y}$ is for the normalisation and can also be written as:

$$p(\mathbf{y}) = \int p(\mathbf{y}|\theta)p(\theta)d\theta$$

Bayes' rule:

$$p(\theta|\mathbf{y}) = \frac{p(\mathbf{y}|\theta)p(\theta)}{p(\mathbf{y})}$$

The marginal distribution of $\mathbf{y}$ is for the normalisation and can also be written as:

$$p(\mathbf{y}) = \int p(\mathbf{y}|\theta)p(\theta)d\theta$$

$$p(\theta|\mathbf{y}) \propto p(\mathbf{y}|\theta)p(\theta)$$

$$p(\boldsymbol{\theta}|\mathbf{y}) \propto p(\mathbf{y}|\boldsymbol{\theta})p(\boldsymbol{\theta})$$

$$p(\boldsymbol{\theta}|\mathbf{y}) \propto p(\mathbf{y}|\boldsymbol{\theta})p(\boldsymbol{\theta})$$

$$\log p(\boldsymbol{\theta}|\mathbf{y}) \propto \log p(\mathbf{y}|\boldsymbol{\theta}) + \log p(\boldsymbol{\theta})$$

$$p(\boldsymbol{\theta}|\mathbf{y}) \propto p(\mathbf{y}|\boldsymbol{\theta})p(\boldsymbol{\theta})$$

$$\log p(\boldsymbol{\theta}|\mathbf{y}) \propto \log p(\mathbf{y}|\boldsymbol{\theta}) + \log p(\boldsymbol{\theta})$$

Instead of Maximum likelihood estimation (MLE), we now do the Maximum A Posterior Estimation (MAP) to get the most likely parameter values

Bayesian Inference

$$p(\mu|y) \propto p(y|\mu) p(\mu)$$

$$p(y|\mu) = \prod_{i=1}^n \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{1}{2} \frac{(y_i - \mu)^2}{\sigma^2}\right) \quad \leftarrow$$

$$p(\mu) = \frac{1}{\sqrt{2\pi\sigma_0^2}} \exp\left(-\frac{1}{2} \frac{(\mu - \mu_0)^2}{\sigma_0^2}\right)$$

$$p(\mu|y) = (2\pi\sigma^2)^{-n/2} \exp\left(-\frac{1}{2\sigma^2} \sum_{i=1}^n (y_i - \mu)^2\right) \\ \cdot (2\pi\sigma_0^2)^{-1/2} \exp\left(-\frac{1}{2\sigma_0^2} (\mu - \mu_0)^2\right)$$

Bayesian Inference

$$\propto \exp \left(- \frac{\sum_{i=1}^n (y_i - \mu)^2}{2\sigma^2} - \frac{(\mu - \mu_0)^2}{2\sigma_0^2} \right)$$

$$\propto \mathcal{N}(\mu_1, \sigma_1^2)$$

$$\sigma_1^2 = \frac{1}{n\sigma^{-2} + \sigma_0^{-2}} \quad \mu_1 = \sigma_1^2 \left(\mu_0 \sigma_0^{-2} + \sum_{i=1}^n y_i \sigma^{-2} \right)$$

Prior distribution -

- is important to reveal any information on the model parameters before observing any data
- can rule out values which are implausible

How to choose a prior?

- To rule out impossible values or distributions

How to choose a prior?

- To rule out impossible values or distributions
- If there is little or no knowledge, uninformative/weakly informative prior distributions can be used i.e. assigning the same probability to all parameter values

How to choose a prior?

- To rule out impossible values or distributions
- If there is little or no knowledge, uninformative/weakly informative prior distributions can be used i.e. assigning the same probability to all parameter values
- Non-informative - Normal distribution with a constant mean and very large variance - $p(\theta)$ is approximately the same for all values of θ

How to choose a prior?

- To rule out impossible values or distributions
- If there is little or no knowledge, uninformative/weakly informative prior distributions can be used i.e. assigning the same probability to all parameter values
- Non-informative - Normal distribution with a constant mean and very large variance - $p(\theta)$ is approximately the same for all values of θ
- Weakly informative - Normal distribution with a constant mean and a large variance (e.g. 50-100) - $p(\theta)$ same for θ which are plausible

How to choose a prior?

- To rule out impossible values or distributions
- If there is little or no knowledge, uninformative/weakly informative prior distributions can be used i.e. assigning the same probability to all parameter values
- Non-informative - Normal distribution with a constant mean and very large variance - $p(\theta)$ is approximately the same for all values of θ
- Weakly informative - Normal distribution with a constant mean and a large variance (e.g. 50-100) - $p(\theta)$ same for θ which are plausible
- Informative - Normal distribution with a constant mean and very small variance

Bayesian Inference

Conjugate priors

Parameter	Prior Distribution
Bernoulli	Beta
Binomial	Beta
Poisson	Gamma
Normal (fixed variance)	Normal for the mean
Normal (fixed mean)	Gamma prior for the inverse variance
Exponential	Gamma
Multinomial	Dirichlet

Conjugacy can help in providing a closed form of the integral, and you do not rely on numerical integration

- There may not be a closed form for the posterior

$$\log p(\boldsymbol{\theta}|\mathbf{y}) \propto \log p(\mathbf{y}|\boldsymbol{\theta}) + \log p(\boldsymbol{\theta})$$

- Markov Chain Monte Carlo (MCMC) can be used to estimate the posterior
- There are several MCMC methods - No U turn sampler, Metropolis-Hastings

Summary

- Three key terms to remember: Prior, Likelihood and Posterior
- Prior - belief before observing or looking at the data - randomness or uncertainty in parameters $p(\theta)$
- Likelihood - probability that any $\mathbf{y} = (y_1, \dots, y_n)$ will occur for given θ
- Posterior - obtained using Bayes rule - complete probability distribution on θ i.e. $p(\theta|\mathbf{y})$ - analytically or using MCMC

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

Gaussian Processes

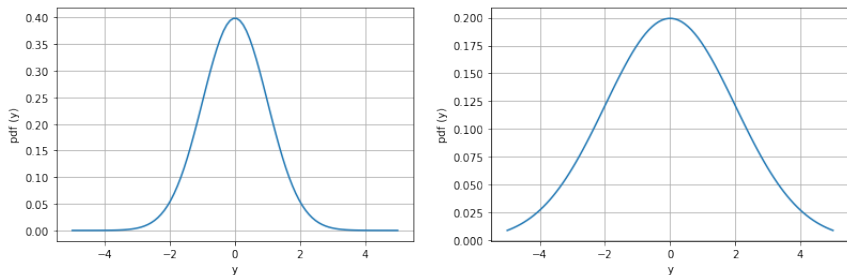


Figure: Two univariate normal distributions with means $(0,0)$ and standard deviation $(1,2)$

Gaussian Processes

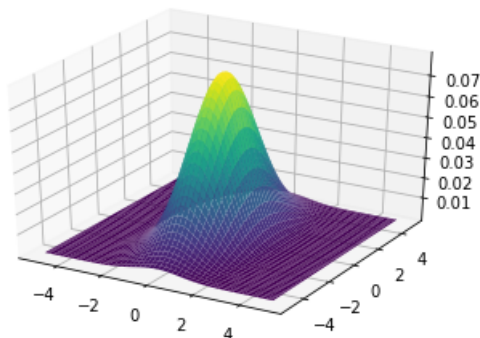


Figure: Joint distribution of two univariate normal distributions is also Normal with means $(0,0)$ and covariance matrix $\begin{bmatrix} 1,0 \\ 0,2 \end{bmatrix}$

Gaussian Processes

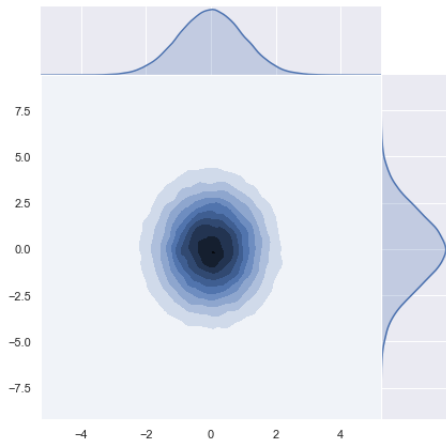


Figure: Joint distribution of two univariate normal distributions is also Normal with means $(0,0)$ and covariance matrix $[1,0; 0,2]$

In regression we need a function $y \sim f(\mathbf{x})$, which can be defined as:

$$\begin{aligned} f(\mathbf{x}) &\sim \mathcal{N}(m(\mathbf{x}), k(\mathbf{x}, \mathbf{x}')) \\ &\sim \mathcal{N}(\mathbf{0}, K) \end{aligned}$$

where $m(\mathbf{x})$ is the mean function and $k(\mathbf{x}, \mathbf{x}')$ is the covariance function (or kernel)

For simplicity in mathematical calculations, we will be using $m(\mathbf{x}) = \mathbf{0}$

Let us use squared exponential (SQ) kernel (also known as RBF or Gaussian kernel):

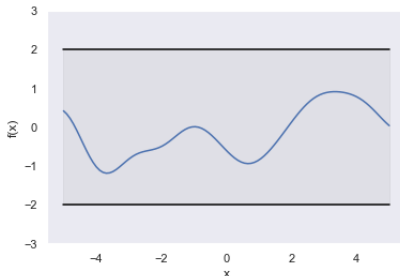
$$\begin{aligned}k(\mathbf{x}, \mathbf{x}') &= \exp(-0.5|\mathbf{x} - \mathbf{x}'|^2) \\ &= \text{cov}(f(\mathbf{x}), f(\mathbf{x}'))\end{aligned}$$

- The squared exponential kernel is infinitely differentiable
- Covariance function between the outputs is used as a function of inputs.
- Specification of covariance function implies a distribution over functions

- Let us draw $f(\mathbf{x})$ with a multivariate Normal distribution using the squared exponential kernel

Gaussian Processes

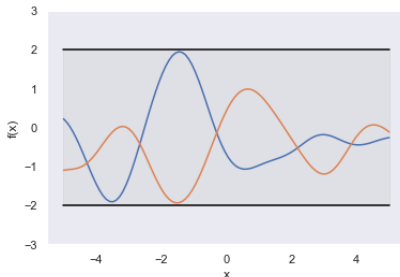
In 1-D: sampling a function randomly using prior distribution over the functions



The prior represent our belief before we observe any data

Gaussian Processes

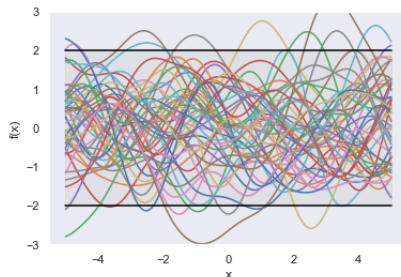
We can have two such functions:



The prior represent our belief before we observe any data

Gaussian Processes

Or we can have infinitely many such functions



The prior represent our belief before we observe any data

Some observations from the prior functions:

- 1 The functions have a *length scale* - can be interpreted as distance you have to move in $\mathbf{x}$ space before $f(\mathbf{x})$ can change significantly

Some observations from the prior functions:

- 1 The functions have a *length scale* - can be interpreted as distance you have to move in $\mathbf{x}$ space before $f(\mathbf{x})$ can change significantly
- 2 Overall variance of process (or joint of all functions) is around one

- To control these two parameters: we can change the formulation of SQ kernel as:

$$k(\mathbf{x}, \mathbf{x}') = \sigma_f^2 \exp\left(-0.5 \frac{|\mathbf{x} - \mathbf{x}'|^2}{l^2}\right)$$

where l is the length scale and σ_f is the amplitude (signal variance) to control the variance

- To control these two parameters: we can change the formulation of SQ kernel as:

$$k(\mathbf{x}, \mathbf{x}') = \sigma_f^2 \exp\left(-0.5 \frac{|\mathbf{x} - \mathbf{x}'|^2}{l^2}\right)$$

where l is the length scale and σ_f is the amplitude (signal variance) to control the variance

- How to set these two parameters - we will come back here
- As for now, assume they are equal to one.

- We are not primarily interested in prior functions
- We are interested in training and prediction

- Given a data set: $(X, \mathbf{y})$, the posterior predictive distribution of new points X^* is also Gaussian:

$$p(\mathbf{y}^* | X^*, X, \mathbf{y}) \sim \mathcal{N}(K(X^* X) K(X, X)^{-1} \mathbf{y}, \\ K(X^*, X^*) - K(X^*, X) K(X, X)^{-1} K(X, X^*))$$

K is covariance matrix in which $K_{ij} = k(\mathbf{x}_i, \mathbf{x}_j)$

- 1 Let us imagine we got some data set for training
- 2 What is the prediction on new samples using Gaussian processes - posterior predictive distribution?

Gaussian Processes

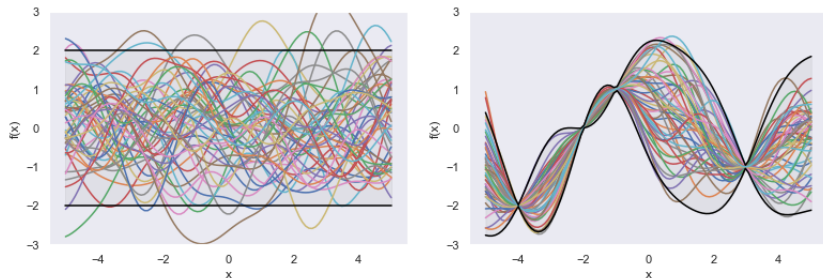


Figure: Prior functions (left) and posterior functions (right) drawn from Gaussian Processes. The black lines represent the mean plus and minus two standard deviations

- For noisy observations,

$$k(\mathbf{x}, \mathbf{x}') = \sigma_f^2 \exp\left(-0.5 \frac{|\mathbf{x} - \mathbf{x}'|^2}{l^2}\right) + \sigma_n^2 \delta_{pq} \quad \text{OR}$$
$$K = K + \sigma_n^2 I,$$

where δ_{pq} is a Kronecker delta which is one iff $p = q$ and zero otherwise and σ_n^2 is the variance of noise which is assumed to be same for same data points (homoscedastic).

Summary so far:

- A stochastic process is a generalisation of a probability distribution to functions
- These functions are Gaussian
- Combination of prior over functions and the data leads to posterior over functions
- For prior, we need to find the covariance function

Summary so far:

- We do not need to worry about basis functions - as they are inbuilt in the covariance function
- Learning in Gaussian process is finding the covariance function
- Covariance function - length scale, amplitude (noise free data) and noise variance (noisy data)

Learning in Gaussian Processes - Finding the parameters of the covariance function

$$k(\mathbf{x}, \mathbf{x}') = \sigma_f^2 \exp\left(-0.5 \frac{|\mathbf{x} - \mathbf{x}'|^2}{l^2}\right) + \sigma_n^2 \delta_{pq},$$

where δ_{pq} is a Kronecker delta which is one iff $p = q$ and zero otherwise.

- We need to maximise the marginal likelihood to get the covariance function parameters

Before we maximise marginal likelihood, let us find the importance of different parameters in the covariance function

Gaussian Processes

Effect of length scale (l):

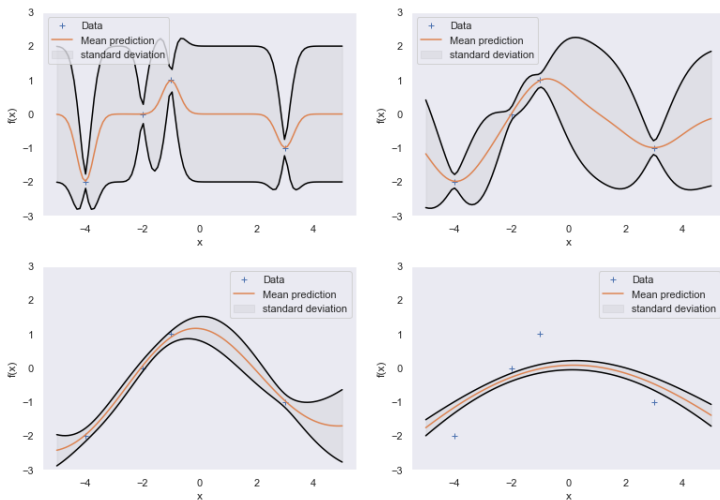


Figure: Effect of length scale, $l = 0.3, 1, 3$ and 10 . $\sigma_f = 1$ and $\sigma_n = 0.1$

Gaussian Processes

Effect of amplitude (σ_f)

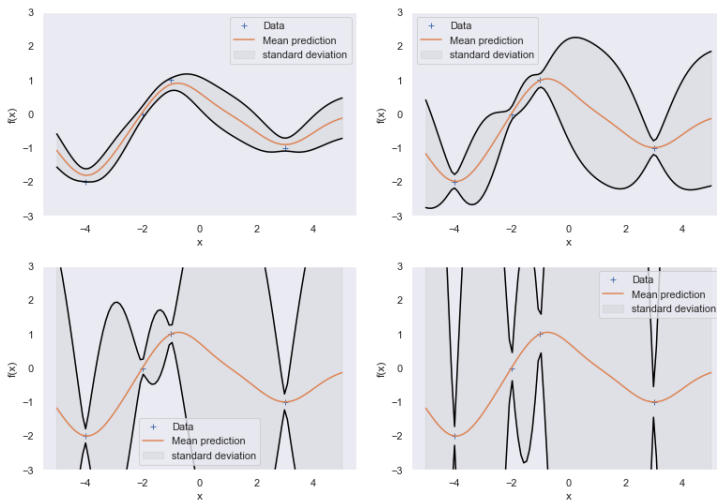


Figure: Effect of amplitude, $\sigma_f = 0.3, 1, 3$ and 10 . $l = 1$ and $\sigma_n = 0.1$

Gaussian Processes

Effect of noise (σ_n)

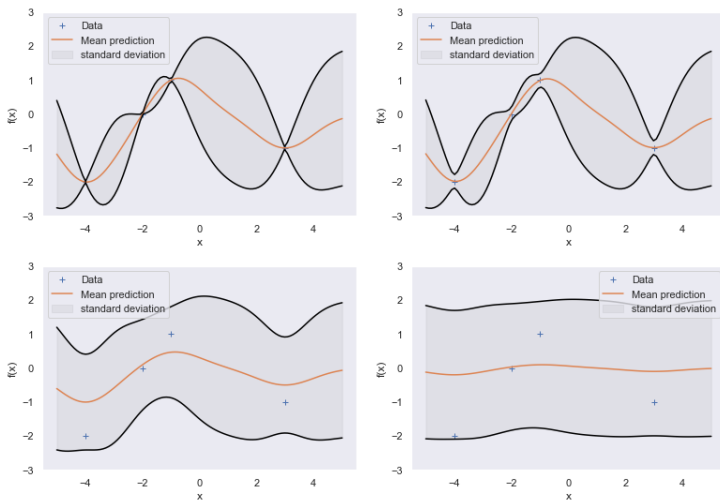


Figure: Effect of amplitude, $\sigma_n = 0, 0.1, 1$ and 3 . $l = 1$ and $\sigma_f = 1$

Estimating the covariance function parameters by maximising the marginal likelihood:

$$p(\mathbf{y}|X) = \int p(\mathbf{y}|\mathbf{f}, X)p(\mathbf{f}|X)d\mathbf{f}$$

It is the marginalisation over functions $\mathbf{y} = f(X)$. As the functions are Gaussian:

$$\log p(\mathbf{y}|X) = -0.5\mathbf{y}^T K^{-1}\mathbf{y} - 0.5|K| - \frac{n}{2} \log 2\pi$$

$$l, \sigma_f, \sigma_n = \underset{l, \sigma_f, \sigma_n}{\operatorname{argmax}} \text{ Marginal Log-Likelihood}$$

After training the model, we are interested in prediction:

$$\mathbf{f}^*|X^*, X, \mathbf{f} \sim \mathcal{N}\left(K(X^*, X)K(X, X)^{-1}\mathbf{f},\right. \\ \left.K(X^*, X^*) - K(X^*, X)K(X, X)^{-1}K(X, X^*)\right)$$

Summary

- **Bayesian Optimisation** - needs uncertainty in prediction in addition to point prediction
- **Gaussian processes** - prior over functions, provide uncertainty in predictions
- **Bayesian models** such as Bayesian NN can also be used in BO

