

Gaussian Processes

* Motivation

1) $f(x) = \sum_{m=1}^M a_m \phi_m(x)$ in Bayesian framework

$$M \rightarrow \infty \Rightarrow GP$$

2) Bayesian Neural Network

$$p(\vec{w} | D) \propto p(D | \vec{w}) p(\vec{w})$$

↳ 1 hidden-layer
with H units

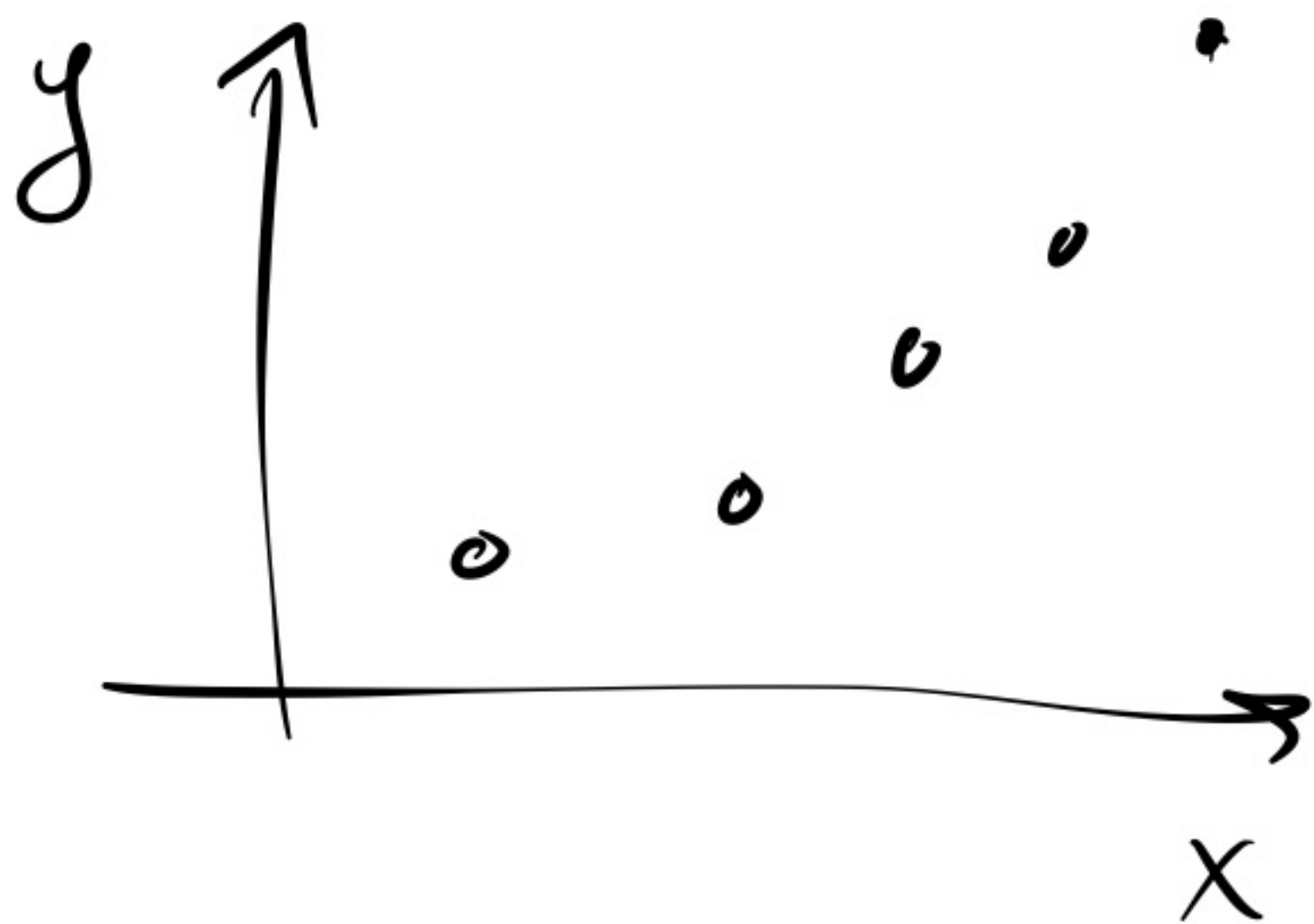
$H \rightarrow \infty \Rightarrow$ distribution of the output
is GP

(Neal '96)

GP flexibility, expressiveness

↳ non-parametric

the complexity grows
with size of
dataset



Parametric approach

1 - assume $p(y)$

e.g.

$$p(y_i) = \mathcal{N}(y_i | f(x_i, \vec{w}), \sigma^2)$$

2 - $f(x_i, \vec{w})$ is given by a ML model

Non parametric

(previous lecture) $\Rightarrow$ relax the 1st step

today $\Rightarrow$ relax 2 step

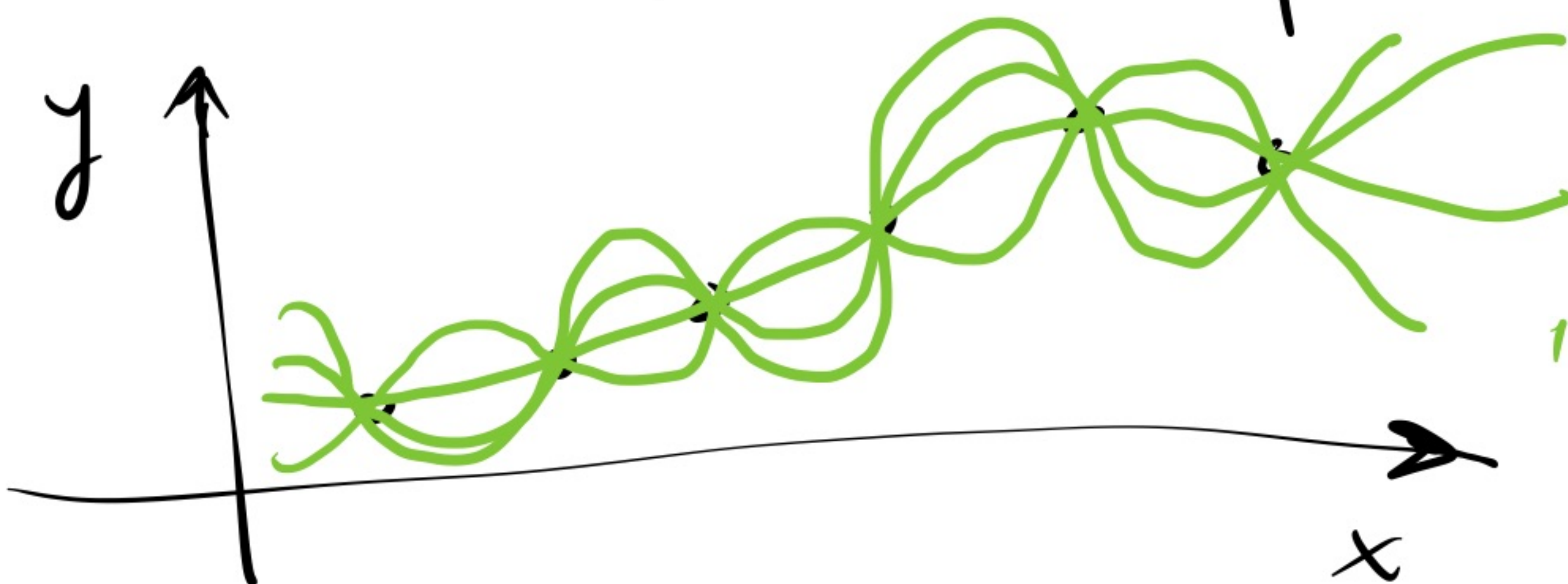
GP:

- Do assume a particular distrib. of your observable $\Rightarrow$ Gaussian

- $f(x, \vec{w}) \Rightarrow$ not fixed

knowing only values $f_i = f(x_i)$

of some finite # of points $\{x_i\}$



$f_i = y_i$

infinite possible

z is random variable

$f(x)$ as a random function

Task define distrib over functions

mean function

$$m(\vec{x}) = \mathbb{E}[f(\vec{x})]$$

covariance of the function

$$\begin{aligned} \mathbb{E}[(f(x) - m(x))(f(x') - m(x')))] \\ \equiv k(x, x') \end{aligned}$$

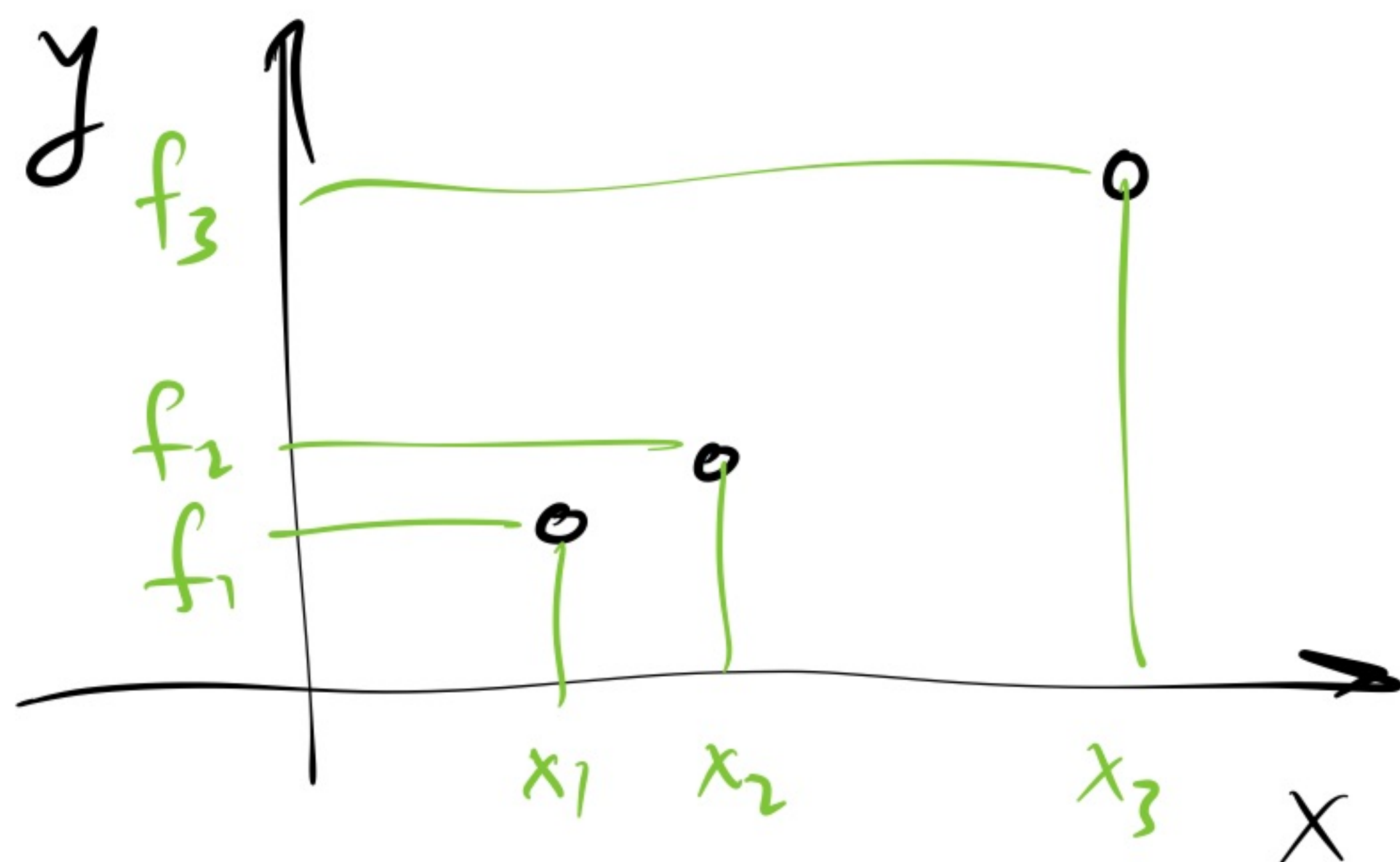
$f(x)$ is a GP

$$f(x) \sim \text{GP}(m(x), k(x, x'))$$

DEF: A GP is a collection of variables,
any finite # of which have a
joint Gaussian distrib.

$\{f_i\}$

Our task : estimate generating function of data



Assume function is smooth



cov. matrix should reflect some measure of similarity

x' close x $\left\{ \begin{array}{l} \Delta \\ f(x') \text{ close } f(x) \end{array} \right\}$
to be correlated

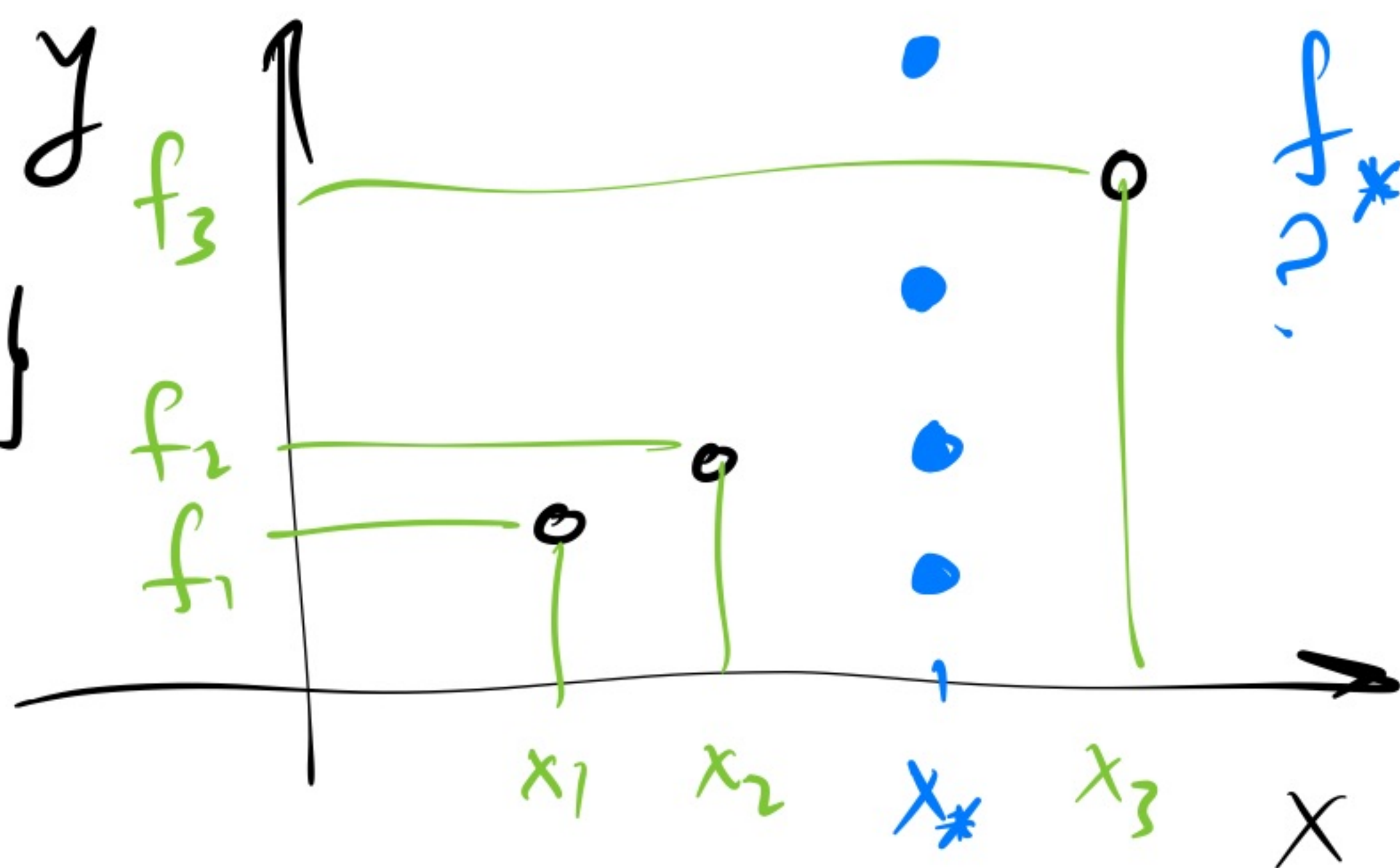
$$\begin{bmatrix} f_1 \\ f_2 \\ f_3 \end{bmatrix} \sim N \left(\begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix}, \begin{bmatrix} 1 & 0.7 & 0.2 \\ 0.7 & 1 & 0.5 \\ 0.2 & 0.5 & 1 \end{bmatrix} \right)$$

kernel function

$$k_{ij} = e^{-\lambda |x_i - x_j|^2} = \begin{cases} 0 & |x_i - x_j| \rightarrow \infty \\ 1 & x_i = x_j \end{cases}$$

Let's evaluate f_*

Given $\vec{f} = \{f_1, f_2, f_3\}$



$$\begin{bmatrix} f_1 \\ f_2 \\ f_3 \\ f_* \end{bmatrix} \sim \mathcal{N} \left(\begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \end{bmatrix}, \begin{bmatrix} \mathbf{K}_{3 \times 3} & \vec{k}_*^T \\ \vec{k}_* & k_{**} \end{bmatrix} \right)$$

$\begin{matrix} 3 \times 3 & 1 \times 3 \\ 3 \times 1 & 1 \times 1 \end{matrix}$

Aim $p(f_* | \vec{f})$

Property

$$\vec{f} = \{f_1, f_2, \dots, f_n\}$$

$$p(\vec{f}) = \mathcal{N}(\vec{f} | \vec{\mu}, \Sigma)$$

$$p(\underbrace{f_1, \dots, f_k}_{\vec{f}_1} | \underbrace{f_{k+1}, \dots, f_n}_{\vec{f}_2}) \Rightarrow \text{Gaussian}$$

with $\vec{f}_1$ mean $\mu_{1/2}$

Cov $\Sigma_{1/2}$

$$p(\vec{f}) = N\left(\begin{bmatrix} \vec{\mu}_1 \\ \vec{\mu}_2 \end{bmatrix}, \begin{bmatrix} \Sigma_{11} & \Sigma_{12} \\ \Sigma_{21} & \Sigma_{22} \end{bmatrix}\right)$$

$$\left[\begin{aligned} \mu_{1/2} &= \vec{\mu}_1 + \Sigma_{12} \Sigma_{22}^{-1} (\vec{f}_2 - \vec{\mu}_2) \\ \Sigma_{1/2} &= \Sigma_{11} - \Sigma_{12} \Sigma_{22}^{-1} \Sigma_{21} \end{aligned} \right] \quad (I)$$

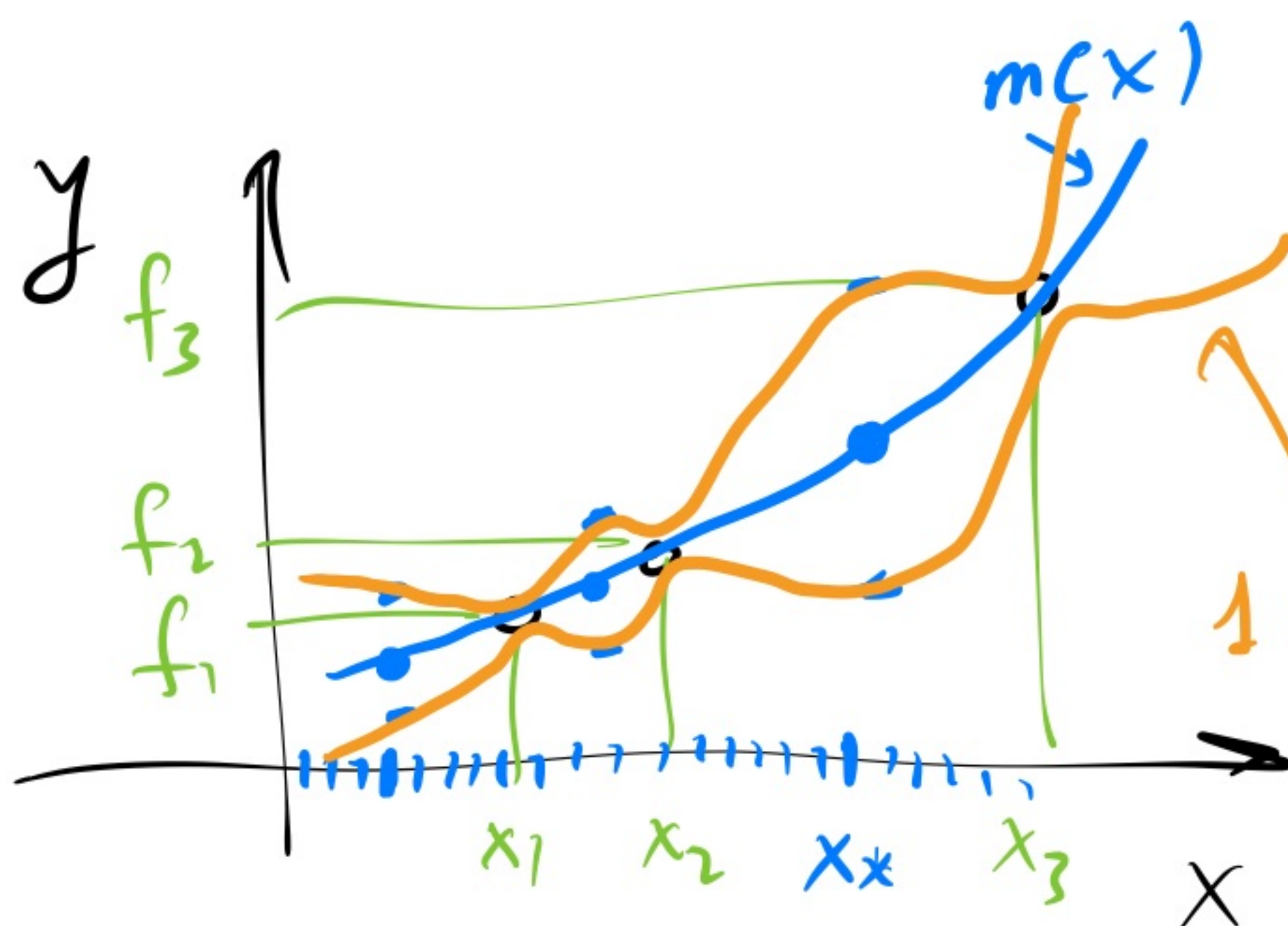
Using (I)

$p(f_* | \vec{f}) \Rightarrow$ Gaussian with

$$\mathbb{E}[f_* | \vec{f}] = \vec{k}_*^T K^{-1} \cdot \vec{f}$$

Variance $\sigma_*^2 = k_{**} - \vec{k}_*^T K^{-1} \cdot \vec{k}_*$

mean and variance of a point in the test set



Go do the same for many x_*

1 std

Summary

- Propose a GP prior
 - create a grid $\{x_i\}_{i=1}^N$
 - compute the matrix $K_{N \times N}$ (covariance)
$$k_{ij} = \exp \left\{ -\frac{1}{2} (x_i - x_j)^2 \right\}$$
 - get values of the function $f_i \equiv f(x_i)$
(sample functions from the prior)

$$\text{Sample } \vec{f} \sim \mathcal{N}(\vec{0}, K) = p(\vec{f})$$

how to do 

$$\text{if } f \sim \mathcal{N}(\mu, \sigma^2)$$

sampling from $\vec{f}$ the same as from

$$f \sim \mu + \sigma \mathcal{N}(0, 1)$$

For multivariate Gaussian

$$\vec{f} \sim \mathcal{N}(\vec{\mu}, K)$$

$$K = L \cdot L^T \Rightarrow \text{Cholesky decomposition}$$

L : real - lower triangular
with positive diag
entries

$$\vec{f} \sim \vec{\mu} + L \cdot N(\vec{0}, \mathbb{I})$$

- Compute $p(f_* | \vec{y})$

Build the Gaussian joint

$$p \left(\begin{matrix} \vec{y} \\ \vec{f}_* \end{matrix} \right) = N \left(\begin{bmatrix} \vec{0}_N \\ \vec{0}_M \end{bmatrix}, \begin{bmatrix} K & K_* \\ K_*^T & K_{**} \end{bmatrix} \right)$$

size N observations $\vec{y}$

predictions for test data size M $\vec{f}_*$

K $N \times N$

K_* $N \times M$

K_*^T $M \times N$

K_{**} $M \times M$

$$\mathbb{E}[\vec{f}_* | \vec{y}] = K_*^T \cdot K^{-1} \cdot \vec{y} = \vec{\mu}_*$$

$$\text{Cov}[\vec{f}_* | \vec{y}] = K_{**} - K_*^T \cdot K^{-1} \cdot K_*$$

$$\Sigma_* = L_* L_*^T$$

- Draw predictions

sample s

$$\vec{f}_*^{(s)} = \vec{\mu}_* + L_* \cdot \vec{\epsilon}^{(s)}$$

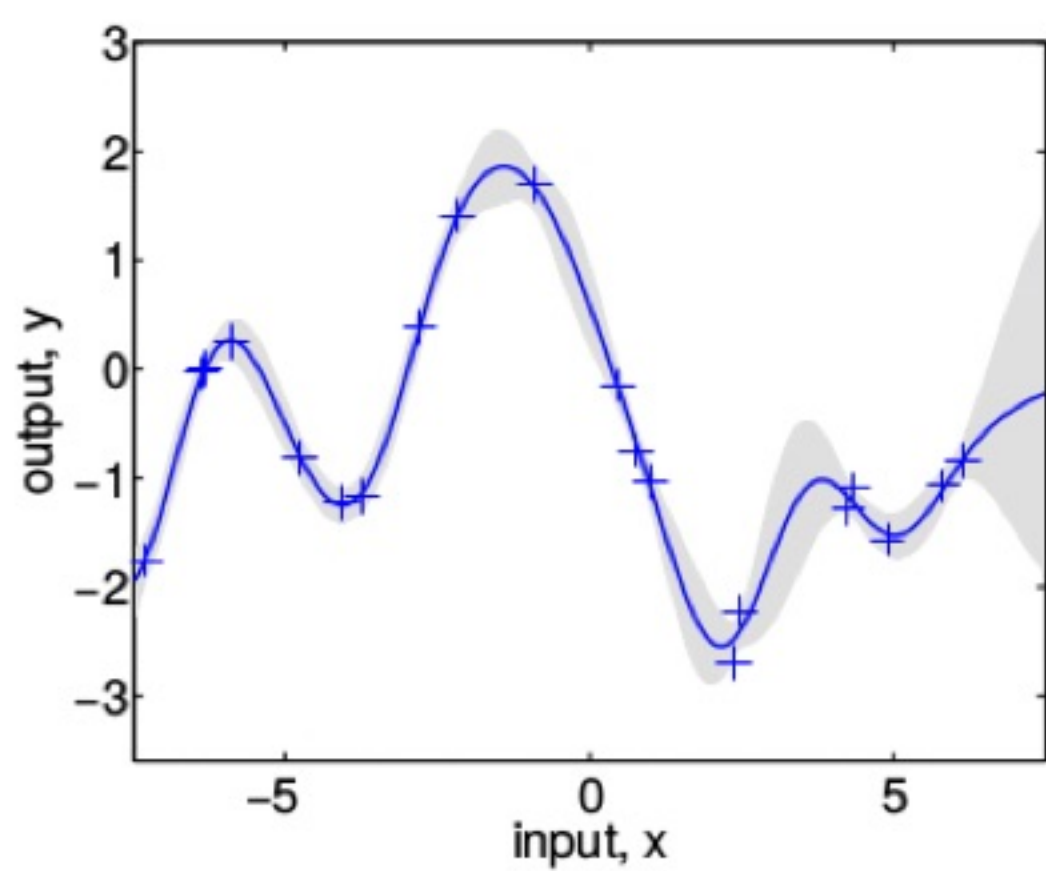
$$\vec{\epsilon}^{(s)} \sim \mathcal{N}(\vec{0}, \mathbb{I})$$

kernel function may have hyperparameters

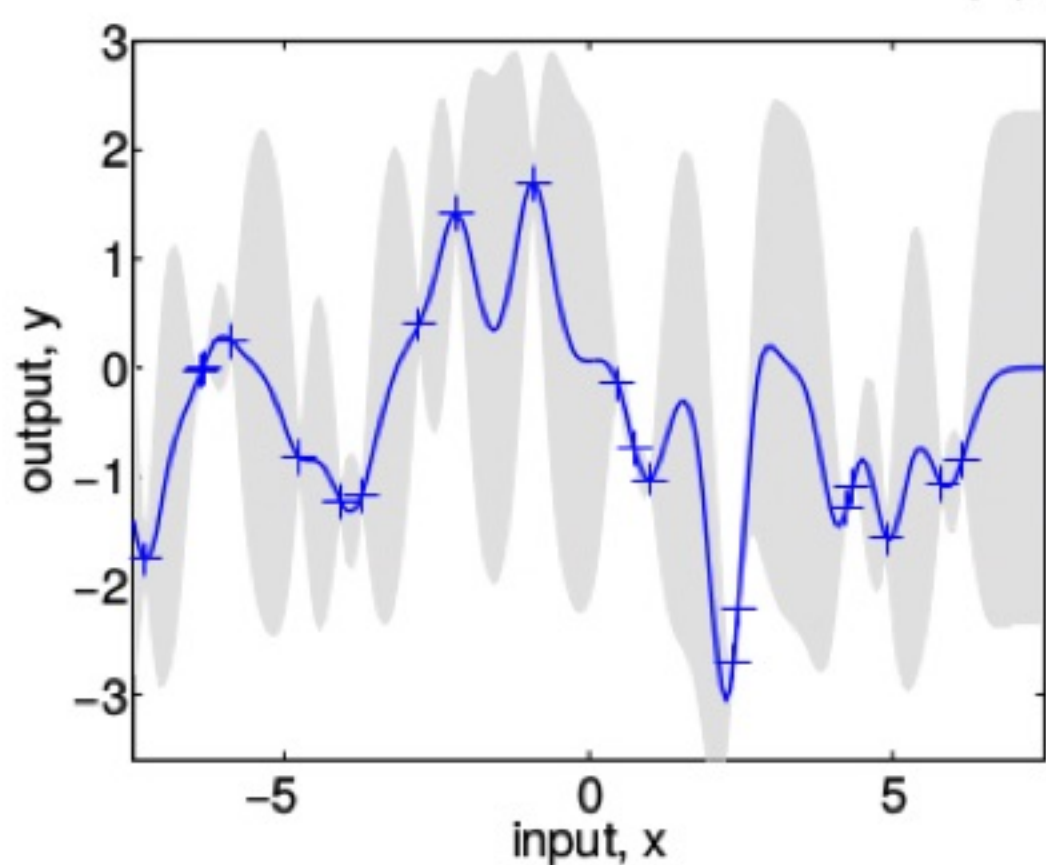
$$k(x, x') = \sigma_f^2 \exp \left\{ -\frac{1}{2\ell^2} (x - x')^2 \right\}$$

a priori should be optimised

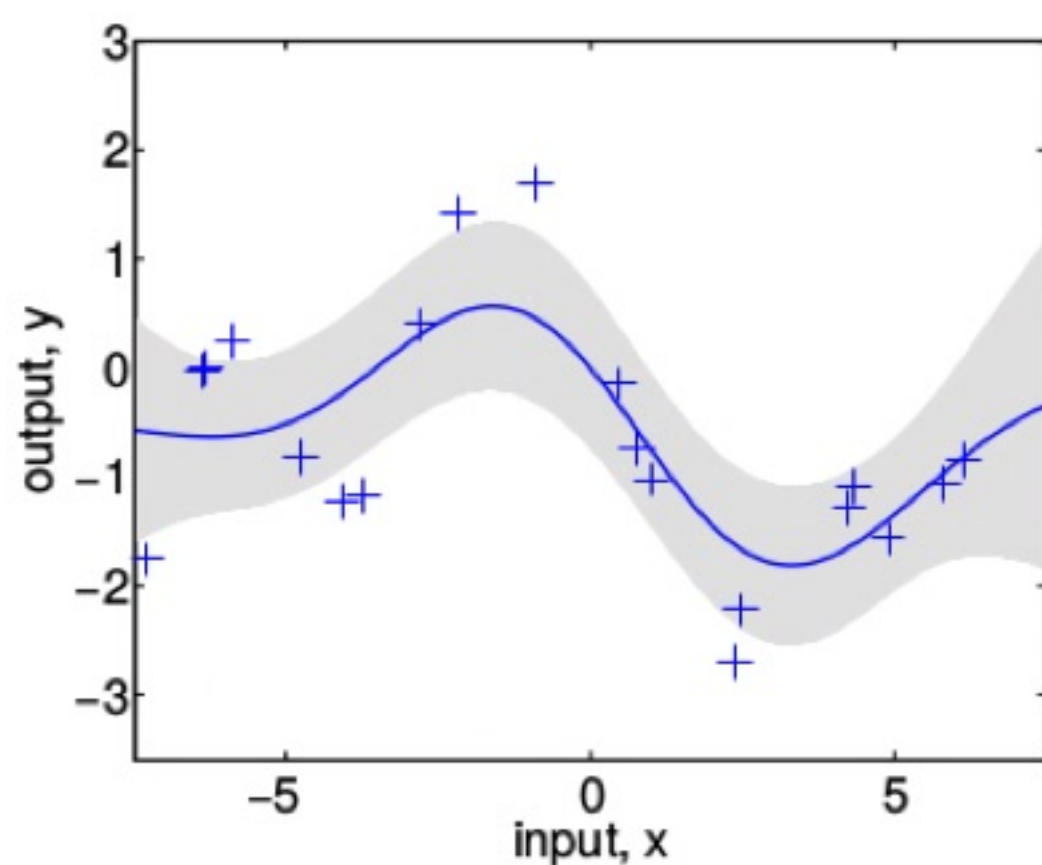
predictions smoother or not depending on



(a), $\ell = 1$



(b), $\ell = 0.3$



(c), $\ell = 3$

(Rasmussen et al.,
"GP for ML")

⇓
should learn the
correct values of
the hyperparameters

you maximise the evidence as a function
of the hyperparameters,

In general noisy data

$$y = f(\vec{x}) + \epsilon$$

$\epsilon \sim N(0, \sigma^2)$

The prior covariance is changed as

$$\vec{f} \sim N(\vec{0}, K + \sigma^2 \mathbb{1})$$

and the predictions

mean: $K_*^T (K + \sigma^2 \mathbb{1})^{-1} \vec{y}$

Cov: $K_{**} - K_*^T (K + \sigma^2 \mathbb{1})^{-1} \cdot K_*$

Classification

$$p(C_k | \vec{x})$$

1-to-k
true labels
} $\vec{y}$

The likelihood

$$P(D | \vec{w}) = \prod_{i=1}^N \prod_{k=1}^K p(C_k | \vec{x}_i, \vec{w})^{1_{i,k}}$$

propose prior $p(\vec{w})$

$\Rightarrow$ Posterior $p(\vec{w} | D)$

The predictive distrib.

$$p(C_k | \vec{x}_*, D) = \int d\vec{w} \underbrace{p(C_k | \vec{x}_*, \vec{w})}_{\text{output of classifier}} p(\vec{w} | D)$$

In GP:

$$\Rightarrow \{f_i\} \in \mathbb{R} \quad -\infty < f_i < \infty$$

$\Downarrow$

$[0, 1]$

“squash” f_i in a bounded function

$$\sigma(f(\vec{x})) \Rightarrow \text{sigmoid } [0, 1]$$

Binary:

$$y \quad p(C=1 | \vec{x}) \quad \text{in binary classif.}$$

Predictive more complicated:

$$P(y_* = c_1 | \vec{x}_*, D)$$

$$= \int \sigma(f_*) P(f_* | \vec{x}_*, D) df_*$$

↳ latent variable

$$P(f_* | \vec{x}_*, D) = \int \underbrace{P(f_* | \vec{x}_*, \vec{f})}_{\text{the conditional as above}} \underbrace{P(\vec{f} | \vec{y})}_{\text{posterior of } \vec{f} \text{ latent functions}} d\vec{f}$$

$$P(\vec{f} | \vec{y}) = \frac{P(\vec{y} | \vec{f}) \underbrace{P(\vec{f} | X)}_{\text{GP prior}}}{P(\vec{y})}$$

↳ intractable

↳ Likelihood (non-Gaussian)

A common choice of likelihood

$$\log p(y_i | f_i) = -\log(1 + e^{-y_i f_i})$$

$$p(y_i | f_i) = \frac{1}{1 + e^{-y_i f_i}}$$

This is reasonable

	$f_i \rightarrow \infty$	$f_i \rightarrow -\infty$
$y_i = +1$	$p(y_i f_i) \rightarrow 1$ <i>high likelihood</i>	$p \rightarrow 0$
$y_i = -1$	$p \rightarrow 0$ <i>low likelihood</i>	$p \rightarrow 1$

For multiclass classif.

instead of 1 latent function

$$\vec{f} = \{f_1, \dots, f_N\}$$

$\hookrightarrow$ # datapoints



1 latent func. per class, K

$$\{\vec{f}_k\} \quad k = 1, \dots, K$$

$$F = \begin{bmatrix} f_{11} & f_{12} & \dots & f_{1K} \\ f_{21} & f_{22} & \dots & f_{2K} \\ \vdots & \vdots & \ddots & \vdots \\ f_{N1} & f_{N2} & \dots & f_{NK} \end{bmatrix}$$

$\Downarrow$ assumed uncorrelated

A popular likelihood

$$P(y_{ik} | \vec{f}_i) = \frac{e^{f_{ik}}}{\sum_{k'} e^{f_{ik'}}}$$

↘ Softmax

Another choice

$$P(y_{ik} | \vec{f}_i) = \prod_{k' \neq k} \theta(f_{ik} - f_{ik'})$$