

10-315: Introduction to Machine Learning

Lecture 23 – Gaussian Processes

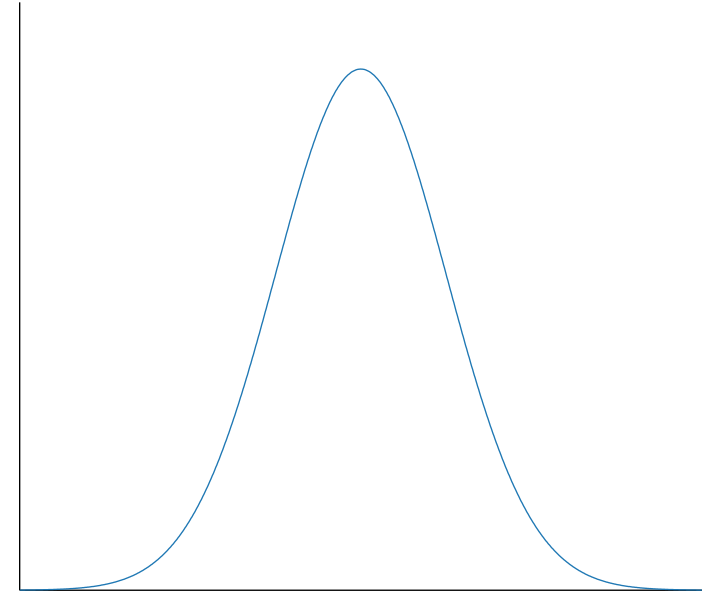
Henry Chai

4/18/22

Gaussians

- (Univariate) Gaussians:

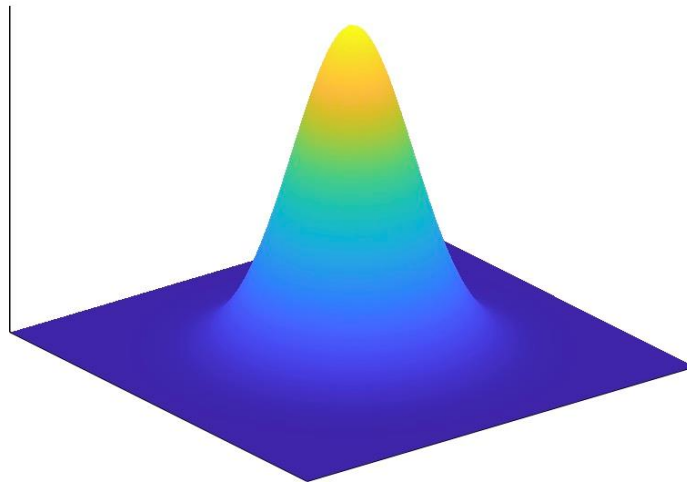
$$x \sim \mathcal{N}(x; \mu = 0, \sigma^2 = 1)$$



- Multivariate Gaussians:

$$\vec{x} = [x_1, \dots, x_p]^T$$

$$\sim \mathcal{N}(\vec{x}; \vec{\mu} = \vec{0}_p, \Sigma = I_p)$$



Some fun facts about Gaussians

- Closure under linear transformations:

$$\text{if } \vec{x} \sim N(\vec{\mu}, \Sigma), \text{ then } C\vec{x} + b \\ \sim N(C\vec{\mu} + b, C\Sigma C^T)$$

- Closure under addition:

$$\text{if } \vec{x} \sim N(\vec{\mu}, \Sigma) \text{ and } \vec{y} \sim N(\vec{m}, S), \\ \text{then } \vec{x} + \vec{y} \sim N(\vec{\mu} + \vec{m}, \Sigma + S)$$

- Closure under conditioning:

$$\text{if } \vec{x} = \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} \sim N\left(\begin{bmatrix} \mu_1 \\ \mu_2 \end{bmatrix}, \begin{bmatrix} \Sigma_{11} & \Sigma_{12} \\ \Sigma_{21} & \Sigma_{22} \end{bmatrix}\right) \\ \text{then } x_1 | x_2 = c \sim N\left(\mu_1 + \Sigma_{12}\Sigma_{22}^{-1}(c - \mu_2), \right. \\ \left. \Sigma_{11} - \Sigma_{12}\Sigma_{22}^{-1}\Sigma_{21}\right)$$

Some old
friends

Gaussian process =

Bayesian linear regression + Kernels

Some old
friends

Gaussian process =

Bayesian linear regression + Kernels

Recall: MAP for Linear Regression

$$A = \begin{bmatrix} 1 & \vec{x}_1 \\ 1 & \vec{x}_2 \\ \vdots & \vdots \\ 1 & \vec{x}_n \end{bmatrix}$$

- If we assume a linear model with additive Gaussian noise

$$\vec{y} = A\vec{\beta} + \vec{\epsilon} \text{ where } \vec{\epsilon} \sim N(\vec{0}_n, \sigma^2 I_n) \rightarrow \vec{y} \sim N(A\vec{\beta}, \sigma^2 I_n)$$

and a Gaussian prior on the weights...

$$\vec{\beta} \sim N\left(\vec{0}_{p+1}, \frac{\sigma^2}{\lambda} I_{p+1}\right) \rightarrow p(\vec{\beta}) \propto \exp\left(-\frac{1}{2\sigma^2} (\lambda \vec{\beta}^T \vec{\beta})\right)$$

- ... then, the MAP of $\vec{\beta}$ is the ridge regression solution!

$$\begin{aligned} \vec{\beta}_{MAP} &= \underset{\vec{\beta}}{\operatorname{argmin}} \left(A\vec{\beta} - \vec{y} \right)^T \left(A\vec{\beta} - \vec{y} \right) + \lambda \vec{\beta}^T \vec{\beta} \\ &= \underline{\left(A^T A + \lambda I_{p+1} \right)^{-1} A^T \vec{y}} \end{aligned}$$

Bayesian Linear Regression

$$A = \begin{bmatrix} 1 & \vec{x}_1 \\ 1 & \vec{x}_2 \\ \vdots & \vdots \\ 1 & \vec{x}_n \end{bmatrix}$$

- Assume a linear model with additive Gaussian noise and a zero-mean Gaussian prior on the weights:

$$\vec{y} = \underline{A\vec{\beta}} + \vec{\epsilon} \text{ where } \vec{\epsilon} \sim N(\underline{\vec{0}_n}, \sigma^2 \underline{I_n}) \text{ and } \underline{\vec{\beta}} \sim N(\underline{\vec{0}_{p+1}}, \underline{\Sigma})$$

- then,

$$\vec{y} \sim N(A\vec{0} + \vec{0}, A\Sigma A^T + \sigma^2 \underline{I_n})$$

Bayesian Linear Regression

$$A = \begin{bmatrix} 1 & \vec{x}_1 \\ 1 & \vec{x}_2 \\ \vdots & \vdots \\ 1 & \vec{x}_n \end{bmatrix}$$

- Assume a linear model with additive Gaussian noise and a zero-mean Gaussian prior on the weights:

$$\vec{y} = A\vec{\beta} + \vec{\epsilon} \text{ where } \vec{\epsilon} \sim N(\vec{0}_n, \sigma^2 I_n) \text{ and } \vec{\beta} \sim N(\vec{0}_{p+1}, \Sigma)$$

- then,

$$\vec{y} \sim N(\vec{0}_n, A\Sigma A^T + \sigma^2 I_n)$$

Bayesian Linear Regression

$$A = \begin{bmatrix} 1 & \vec{x}_1 \\ 1 & \vec{x}_2 \\ \vdots & \vdots \\ 1 & \vec{x}_n \end{bmatrix}$$

- Assume a linear model with additive Gaussian noise and a zero-mean Gaussian prior on the weights:

$$\vec{y} = A\vec{\beta} + \vec{\epsilon} \text{ where } \vec{\epsilon} \sim N(\vec{0}_n, \sigma^2 I_n) \text{ and } \vec{\beta} \sim N(\vec{0}_{p+1}, \Sigma)$$

- then,

$$\begin{bmatrix} \vec{\beta} \\ \vec{y} \end{bmatrix} \sim N \left(\begin{bmatrix} \vec{0}_{p+1} \\ \vec{0}_n \end{bmatrix}, \begin{bmatrix} \Sigma & ??? \\ ??? & A\Sigma A^T + \sigma^2 I_n \end{bmatrix} \right)$$

- Covariance between $\vec{y}$ and $\vec{\beta}$:

$$\begin{aligned} \text{Cov}(\vec{y}, \vec{\beta}) &= \text{Cov}(A\vec{\beta} + \vec{\epsilon}, \vec{\beta}) = \text{Cov}(A\vec{\beta}, \vec{\beta}) \\ &= A \text{Cov}(\vec{\beta}, \vec{\beta}) = A\Sigma \end{aligned}$$

Bayesian Linear Regression

$$A = \begin{bmatrix} 1 & \vec{x}_1 \\ 1 & \vec{x}_2 \\ \vdots & \vdots \\ 1 & \vec{x}_n \end{bmatrix}$$

- Assume a linear model with additive Gaussian noise and a zero-mean Gaussian prior on the weights:

$$\vec{y} = A\vec{\beta} + \vec{\epsilon} \text{ where } \vec{\epsilon} \sim N(\vec{0}_n, \sigma^2 I_n) \text{ and } \vec{\beta} \sim N(\vec{0}_{p+1}, \Sigma)$$

- then,

$$\begin{bmatrix} \vec{\beta} \\ \vec{y} \end{bmatrix} \sim N \left(\begin{bmatrix} \vec{0}_{p+1} \\ \vec{0}_n \end{bmatrix}, \begin{bmatrix} \Sigma & \Sigma A^T \\ A\Sigma & A\Sigma A^T + \sigma^2 I_n \end{bmatrix} \right)$$

- Covariance between $\vec{y}$ and $\vec{\beta}$:

$$\text{Cov}(\vec{y}, \vec{\beta}) = \text{Cov}(A\vec{\beta} + \vec{\epsilon}, \vec{\beta}) = A\text{Cov}(\vec{\beta}, \vec{\beta}) = A\Sigma$$

Bayesian Linear Regression

$$A = \begin{bmatrix} 1 & \vec{x}_1 \\ 1 & \vec{x}_2 \\ \vdots & \vdots \\ 1 & \vec{x}_n \end{bmatrix}$$

- Assume a linear model with additive Gaussian noise and a zero-mean Gaussian prior on the weights:

$$\vec{y} = A\vec{\beta} + \vec{\epsilon} \text{ where } \vec{\epsilon} \sim N(\vec{0}_n, \sigma^2 I_n) \text{ and } \vec{\beta} \sim N(\vec{0}_{p+1}, \Sigma)$$

- then,

$$\vec{\beta} | \vec{y} \sim N(\vec{\mu}_{POST}, \Sigma_{POST})$$

where

$$\vec{\mu}_{POST} = \Sigma A^T (A \Sigma A^T + \sigma^2 I_n)^{-1} \vec{y},$$

$$\Sigma_{POST} = \Sigma - \Sigma A^T (A \Sigma A^T + \sigma^2 I_n)^{-1} A \Sigma$$

Bayesian Linear Regression

$$A = \begin{bmatrix} 1 & \vec{x}_1 \\ 1 & \vec{x}_2 \\ \vdots & \vdots \\ 1 & \vec{x}_n \end{bmatrix}$$

- Assume a linear model with additive Gaussian noise and a zero-mean Gaussian prior on the weights:

$$\vec{y} = A\vec{\beta} + \vec{\epsilon} \text{ where } \vec{\epsilon} \sim N(\vec{0}_n, \sigma^2 I_n) \text{ and } \vec{\beta} \sim N(\vec{0}_{p+1}, \Sigma)$$

- then given a new test data point $\vec{x}^*$, the prediction is

$$y^* | \vec{y} = \vec{x}^{*T} \vec{\beta} | \vec{y} \sim N(\underbrace{\vec{x}^{*T} \vec{\mu}_{POST}}_{\text{mean}}, \underbrace{\vec{x}^{*T} \Sigma_{POST} \vec{x}^*}_{\text{variance}})$$

where

$$\vec{\mu}_{POST} = \Sigma A^T (A \Sigma A^T + \sigma^2 I_n)^{-1} \vec{y},$$

$$\Sigma_{POST} = \Sigma - \Sigma A^T (A \Sigma A^T + \sigma^2 I_n)^{-1} A \Sigma$$

Bayesian Linear Regression

$$A = \begin{bmatrix} 1 & \vec{x}_1 \\ 1 & \vec{x}_2 \\ \vdots & \vdots \\ 1 & \vec{x}_n \end{bmatrix}$$

- Assume a linear model with additive Gaussian noise and a zero-mean Gaussian prior on the weights:

$$\vec{y} = A\vec{\beta} + \vec{\epsilon} \text{ where } \vec{\epsilon} \sim N(\vec{0}_n, \sigma^2 I_n) \text{ and } \vec{\beta} \sim N(\vec{0}_{p+1}, \Sigma)$$

- then given a new test data point $\vec{x}^*$, the prediction is

$$y^* | \vec{y} = \vec{x}^{*T} \vec{\beta} | \vec{y} \sim N(\underline{\vec{\mu}_{PRED}}, \underline{\Sigma_{PRED}})$$

where

$$\vec{\mu}_{PRED} = \vec{x}^{*T} \Sigma A^T (A \Sigma A^T + \sigma^2 I_n)^{-1} \vec{y},$$

$$\Sigma_{PRED} = \vec{x}^{*T} \Sigma \vec{x}^* - \vec{x}^{*T} \Sigma A^T (A \Sigma A^T + \sigma^2 I_n)^{-1} A \Sigma \vec{x}^*$$

Some old
friends

Gaussian process =

Bayesian linear regression + Kernels

Some old
friends

Gaussian process =
Bayesian linear regression + **Kernels**

Bayesian Linear Regression...

$$A = \begin{bmatrix} 1 & \vec{x}_1 \\ 1 & \vec{x}_2 \\ \vdots & \vdots \\ 1 & \vec{x}_n \end{bmatrix}$$

- Assume a linear model with additive Gaussian noise and a zero-mean Gaussian prior on the weights:

$$\vec{y} = A\vec{\beta} + \vec{\epsilon} \text{ where } \vec{\epsilon} \sim N(\vec{0}_n, \sigma^2 I_n) \text{ and } \vec{\beta} \sim N(\vec{0}_{p+1}, \Sigma)$$

- then given a new test data point $\vec{x}^*$, the prediction is

$$y^* | \vec{y} = \vec{x}^{*T} \vec{\beta} | \vec{y} \sim N(\vec{\mu}_{PRED}, \Sigma_{PRED})$$

where

$$\vec{\mu}_{PRED} = \vec{x}^{*T} \Sigma A^T (A \Sigma A^T + \sigma^2 I_n)^{-1} \vec{y},$$

$$\Sigma_{POST} = \vec{x}^{*T} \Sigma \vec{x}^* - \vec{x}^{*T} \Sigma A^T (A \Sigma A^T + \sigma^2 I_n)^{-1} A \Sigma \vec{x}^*$$

Bayesian Linear Regression can be kernelized!

$$\underline{\Phi} = \begin{bmatrix} 1 & \underline{\phi(\vec{x}_1)} \\ 1 & \underline{\phi(\vec{x}_2)} \\ \vdots & \vdots \\ 1 & \underline{\phi(\vec{x}_n)} \end{bmatrix}$$

- Assume a linear model with additive Gaussian noise and a zero-mean Gaussian prior on the weights:

$$\vec{y} = \Phi \vec{\beta} + \vec{\epsilon} \text{ where } \vec{\epsilon} \sim N(\vec{0}_n, \sigma^2 I_n) \text{ and } \vec{\beta} \sim N(\vec{0}_{p'+1}, \Sigma)$$

- then given a new test data point $\underline{\vec{x}^*}$, the prediction is

$$y^* | \vec{y} = \underline{\phi(\vec{x}^*)}^T \vec{\beta} | \vec{y} \sim N(\vec{\mu}_{PRED}, \Sigma_{PRED})$$

where

$$\vec{\mu}_{PRED} = \underline{\phi(\vec{x}^*)}^T \underline{\Sigma} \underline{\Phi}^T (\Phi \Sigma \Phi^T + \sigma^2 I_n)^{-1} \vec{y},$$

Σ_{POST}

$$= \phi(\vec{x}^*)^T \Sigma \phi(\vec{x}^*) - \phi(\vec{x}^*)^T \Sigma \Phi^T (\Phi \Sigma \Phi^T + \sigma^2 I_n)^{-1} \Phi \Sigma \phi(\vec{x}^*)$$

Bayesian Linear Regression can be kernelized!

$$\Phi = \begin{bmatrix} 1 & \phi(\vec{x}_1) \\ 1 & \phi(\vec{x}_2) \\ \vdots & \vdots \\ 1 & \phi(\vec{x}_n) \end{bmatrix}$$

- Assume a linear model with additive Gaussian noise and a zero-mean Gaussian prior on the weights:

$$\vec{y} = \Phi \vec{\beta} + \vec{\epsilon} \text{ where } \vec{\epsilon} \sim N(\vec{0}_n, \sigma^2 I_n) \text{ and } \vec{\beta} \sim N(\vec{0}_{p'+1}, \Sigma)$$

- then given a new test data point $\vec{x}^*$, the prediction is

$$y^* | \vec{y} = \phi(\vec{x}^*)^T \vec{\beta} | \vec{y} \sim N(\vec{\mu}_{PRED}, \Sigma_{PRED})$$

where

$$\vec{\mu}_{PRED} = \phi(\vec{x}^*)^T \Sigma \Phi^T (\Phi \Sigma \Phi^T + \sigma^2 I_n)^{-1} \vec{y},$$

$$\Sigma_{POST}$$

$$= \phi(\vec{x}^*)^T \Sigma \phi(\vec{x}^*) - \phi(\vec{x}^*)^T \Sigma \Phi^T (\Phi \Sigma \Phi^T + \sigma^2 I_n)^{-1} \Phi \Sigma \phi(\vec{x}^*)$$

- Define the kernel function to be

$$\underline{K(\vec{x}, \vec{x}')} = \phi(\vec{x})^T \Sigma \phi(\vec{x}')$$

$$= UU^T$$

Σ is p.s.d.

Training data

$$\mathcal{D} = \{(\vec{x}_i, y_i)\}_{i=1}^n$$

Bayesian
Linear
Regression can
be kernelized!

$$K(A, A) \in \mathbb{R}^{n \times n}$$

$$K(\vec{x}_1, \vec{x}_1) \quad K(\vec{x}_1, \vec{x}_2) \quad \dots$$

$$K(\vec{x}_2, \vec{x}_1) \quad K(\vec{x}_2, \vec{x}_2) \quad \dots$$

$$\vdots \quad \vdots \quad \ddots$$

- Assume a linear model with additive Gaussian noise and a zero-mean Gaussian prior on the weights:

$$\vec{y} = \Phi \vec{\beta} + \vec{\epsilon} \text{ where } \vec{\epsilon} \sim N(\vec{0}_n, \sigma^2 I_n) \text{ and } \vec{\beta} \sim N(\vec{0}_{p'+1}, \Sigma)$$

- then given a new test data point $\vec{x}^*$, the prediction is

$$\underline{y^* | \vec{y}} \sim N(\underline{\vec{\mu}_{PRED}}, \underline{\Sigma_{PRED}})$$

where

$$\underline{\vec{\mu}_{PRED}} = \underline{K(\vec{x}^*, A)} (K(A, A) + \sigma^2 I_n)^{-1} \vec{y},$$

$$\underline{\Sigma_{POST}} = \underline{K(\vec{x}^*, \vec{x}^*)} - \underline{K(\vec{x}^*, A)} (K(A, A) + \sigma^2 I_n)^{-1} K(A, \vec{x}^*)$$

$$K(\vec{x}^*, A) = \begin{bmatrix} K(\vec{x}^*, \vec{x}_1) \\ K(\vec{x}^*, \vec{x}_2) \\ \vdots \\ K(\vec{x}^*, \vec{x}_n) \end{bmatrix}$$

- Define the kernel function to be

$$K(\vec{x}, \vec{x}') = \phi(\vec{x})^T \Sigma \phi(\vec{x}')$$

Wait, what happened to the weights?

- Assume a linear model with additive Gaussian noise and a zero-mean Gaussian prior on the weights:

$$\vec{y} = \Phi \vec{\beta} + \vec{\epsilon} \text{ where } \vec{\epsilon} \sim N(\vec{0}_n, \sigma^2 I_n) \text{ and } \vec{\beta} \sim N(\vec{0}_{p'+1}, \Sigma)$$

- then given a new test data point $\vec{x}^*$, the prediction is

$$y^* | \vec{y} \sim N(\vec{\mu}_{PRED}, \Sigma_{PRED})$$

where

$$\vec{\mu}_{PRED} = K(\vec{x}^*, A) (K(A, A) + \sigma^2 I_n)^{-1} \vec{y},$$

$$\Sigma_{POST} = K(\vec{x}^*, \vec{x}^*) - K(\vec{x}^*, A) (K(A, A) + \sigma^2 I_n)^{-1} K(A, \vec{x}^*)$$

- Define the kernel function to be

$$K(\vec{x}, \vec{x}') = \phi(\vec{x})^T \Sigma \phi(\vec{x}')$$

Some old
friends

Gaussian process =
Bayesian linear regression + Kernels

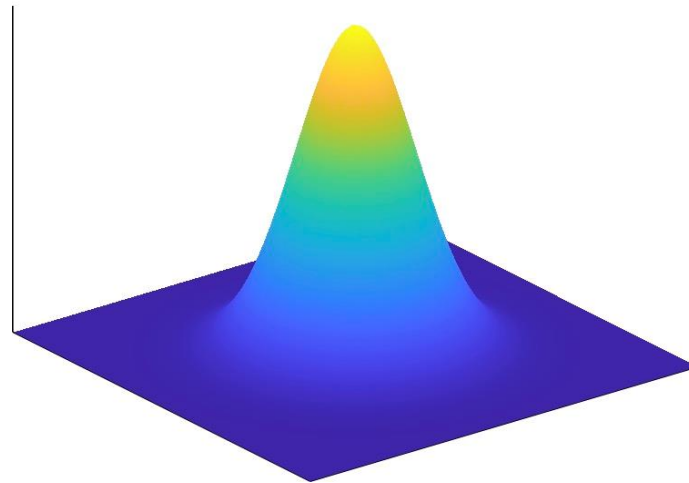
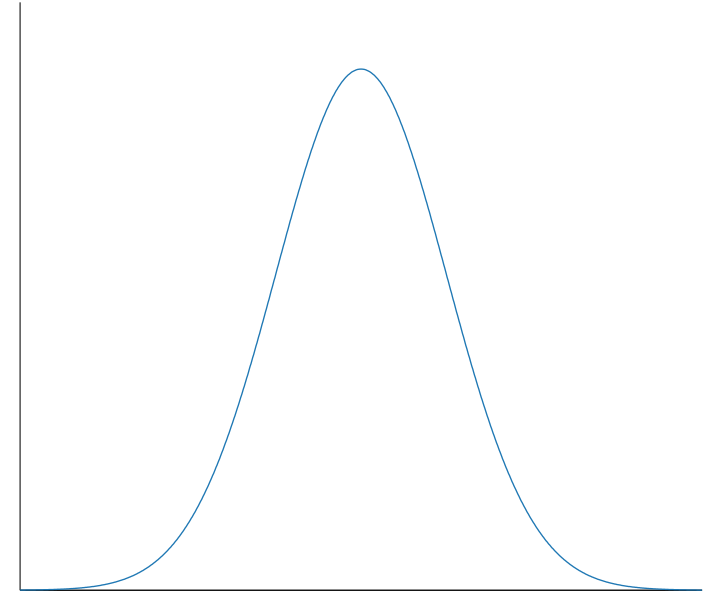
A new
perspective

Gaussian process =
The extension of a Gaussian
distribution to functions

Gaussians

- (Univariate) Gaussians:

$$x \sim \mathcal{N}(x; \mu = 0, \sigma^2 = 1)$$



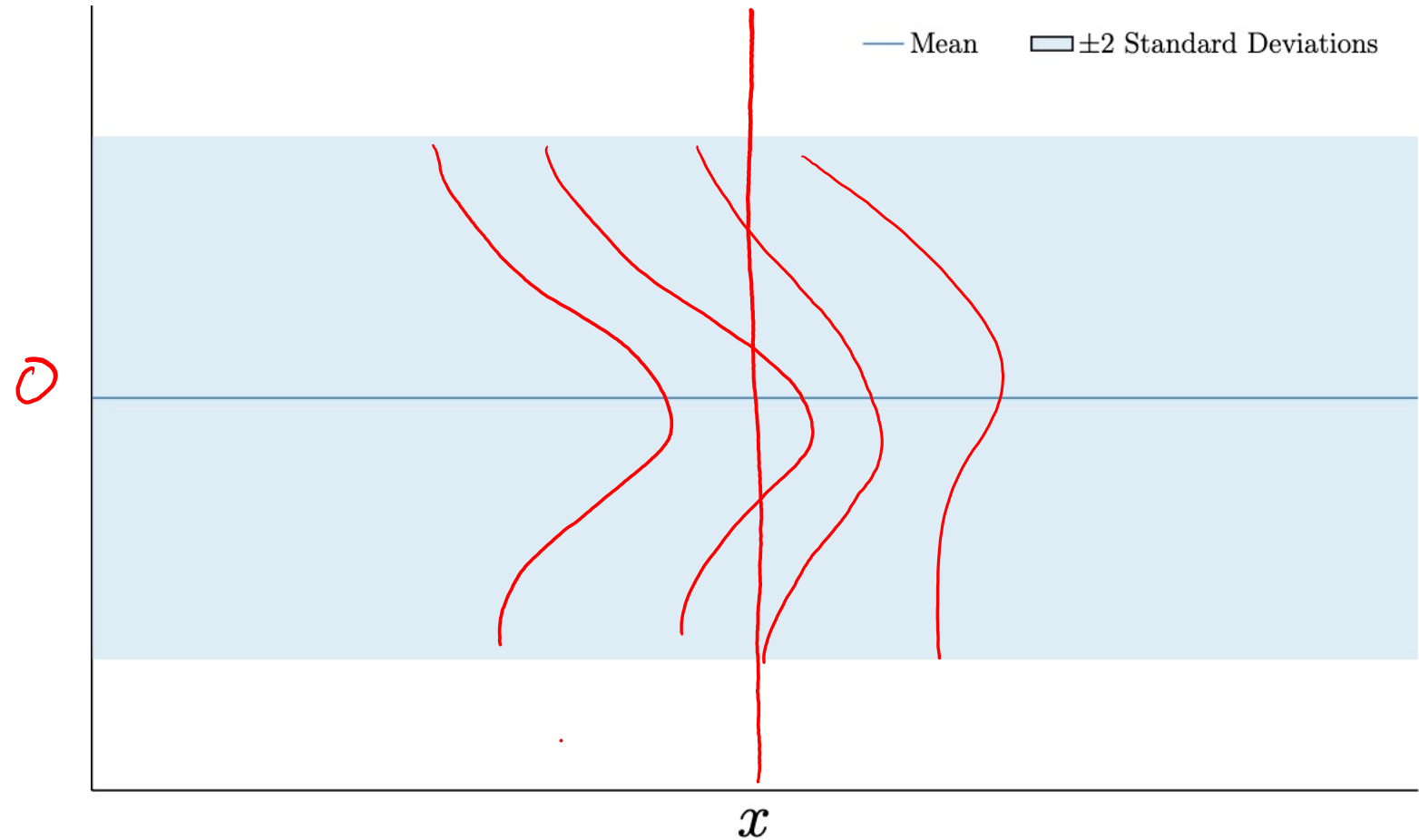
- Multivariate Gaussians:

$$\vec{x} = [x_1, \dots, x_p]^T$$

$$\sim \mathcal{N}(\vec{x}; \vec{\mu} = \vec{0}_p, \Sigma = I_p)$$

Gaussian Process (GP)

$$f: \mathbb{R}^p \mapsto \mathbb{R} \sim \mathcal{GP}(f; \mu(x), \Sigma(x, x'))$$



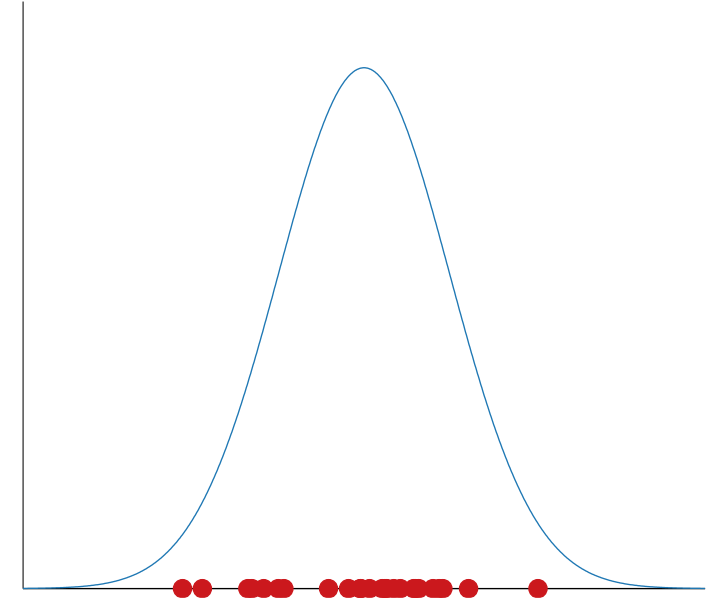
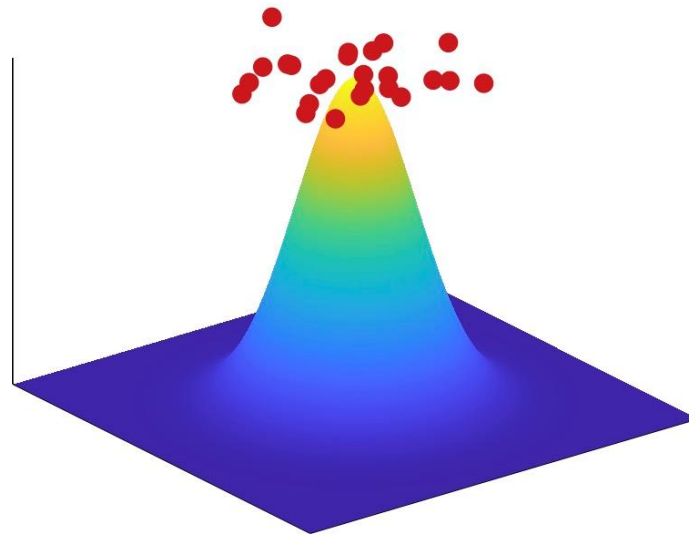
$$f \sim \mathcal{GP}(\mu, \Sigma) \rightarrow f(x) \sim \mathcal{N}(\mu(x), \Sigma(x, x))$$

Squared exponential

Gaussians

- (Univariate) Gaussians:

$$x \sim \mathcal{N}(x; \mu = 0, \sigma^2 = 1)$$



- Multivariate Gaussians:

$$\vec{x} = [x_1, \dots, x_p]^T$$

$$\sim \mathcal{N}(\vec{x}; \vec{\mu} = \vec{0}_p, \Sigma = I_p)$$

$$\frac{\partial f}{\partial x} \sim \mathcal{GP}\left(\frac{\partial \mu}{\partial x}, \frac{\partial^2}{\partial x \partial x'}\right)$$

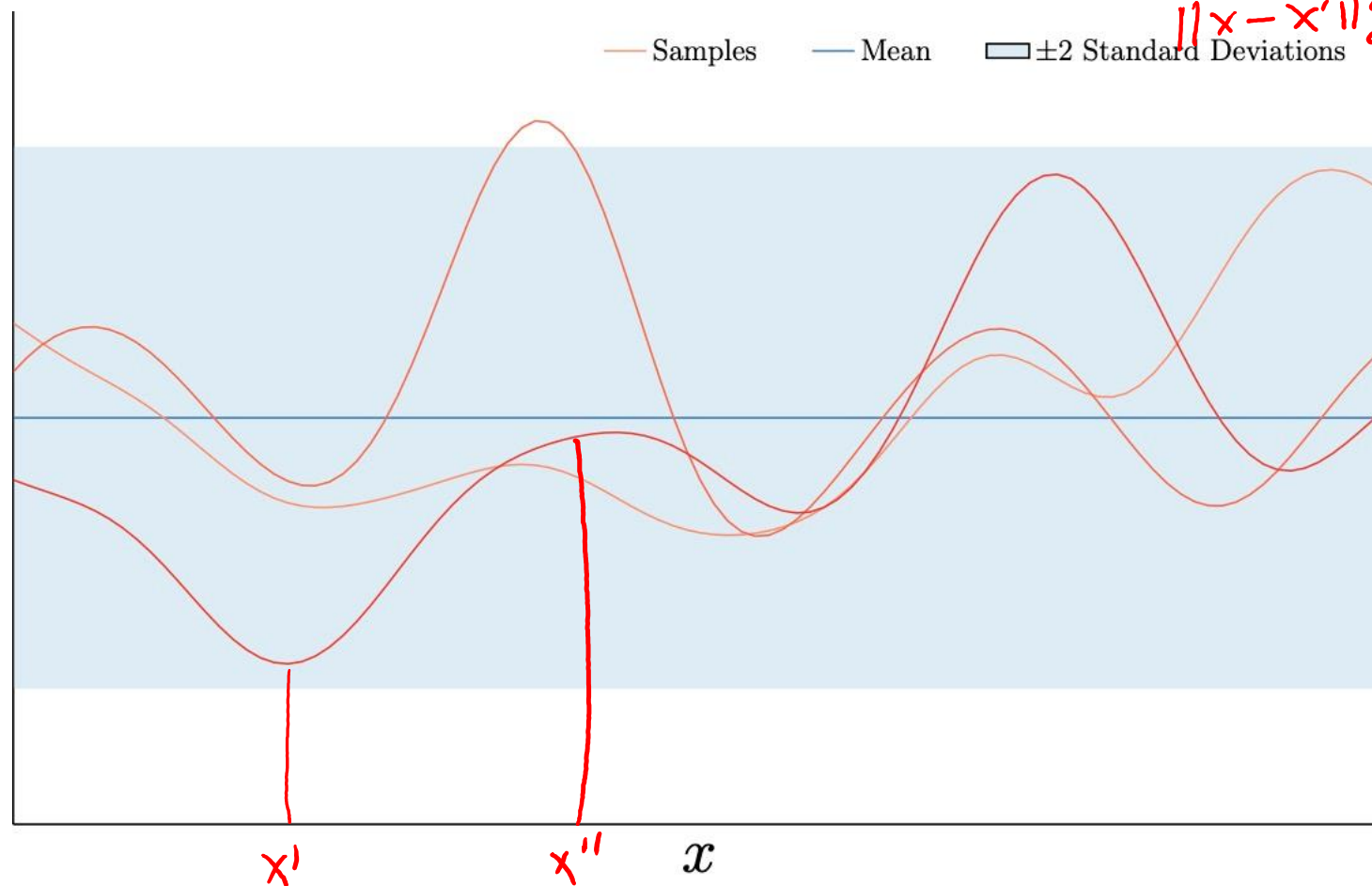
Gaussian Process (GP)

$\Sigma(x', x'')$
 = covariance of
 $f(x')$ and $f(x'')$

$$f: \mathbb{R}^p \mapsto \mathbb{R} \sim \mathcal{GP}(f; \mu(x) = 0, \Sigma(x, x') = \exp(-(x - x')^2))$$

if x and x'
 are far apart ≈ 0

$\frac{1}{\sqrt{2}} \|x - x'\|_2$

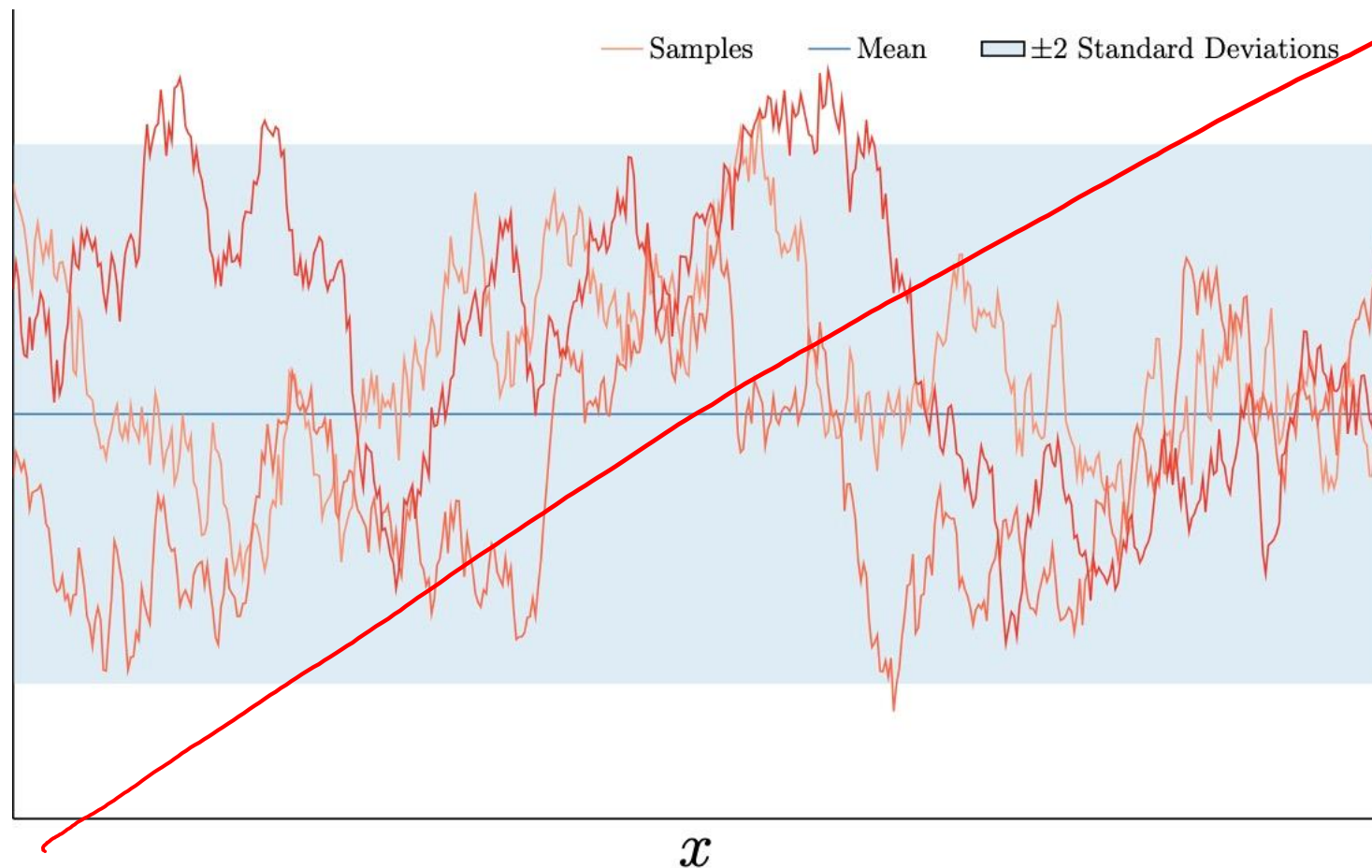


$$f \sim \mathcal{GP}(\mu, \Sigma) \rightarrow f(x) \sim \mathcal{N}(\mu(x), \Sigma(x, x))$$

$e^{\log(f)} \sim \mathcal{GP}$
 $f \geq 0$

Gaussian Process (GP)

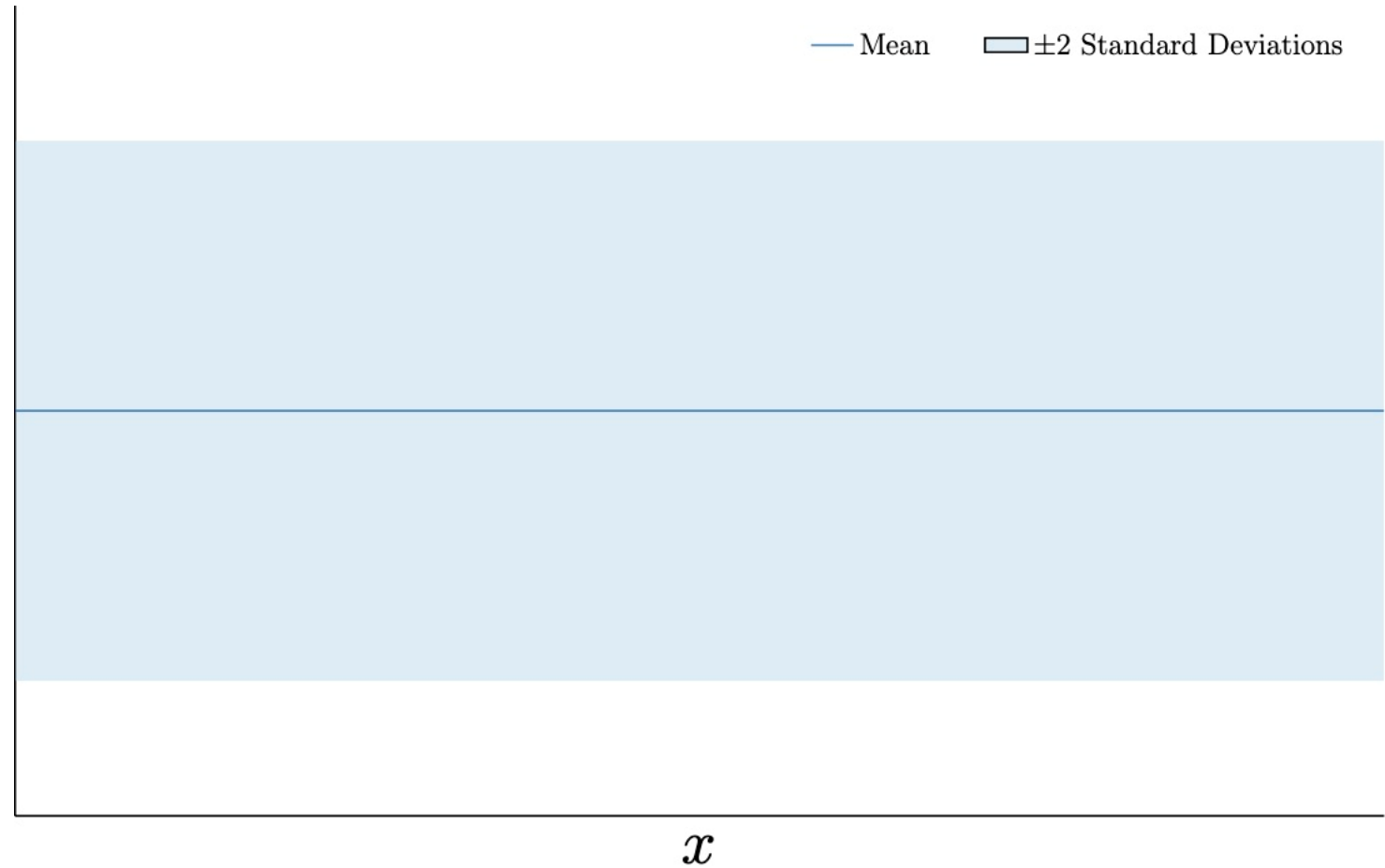
$$f: \mathbb{R}^p \mapsto \mathbb{R} \sim \mathcal{GP}(f; \mu(x) = 0, \Sigma(x, x') = \underbrace{\exp(-\|x - x'\|)}_{\sin(\|x - x'\|)})$$



$$f \sim \mathcal{GP}(\mu, \Sigma) \rightarrow f(x) \sim \mathcal{N}(\mu(x), \Sigma(x, x))$$

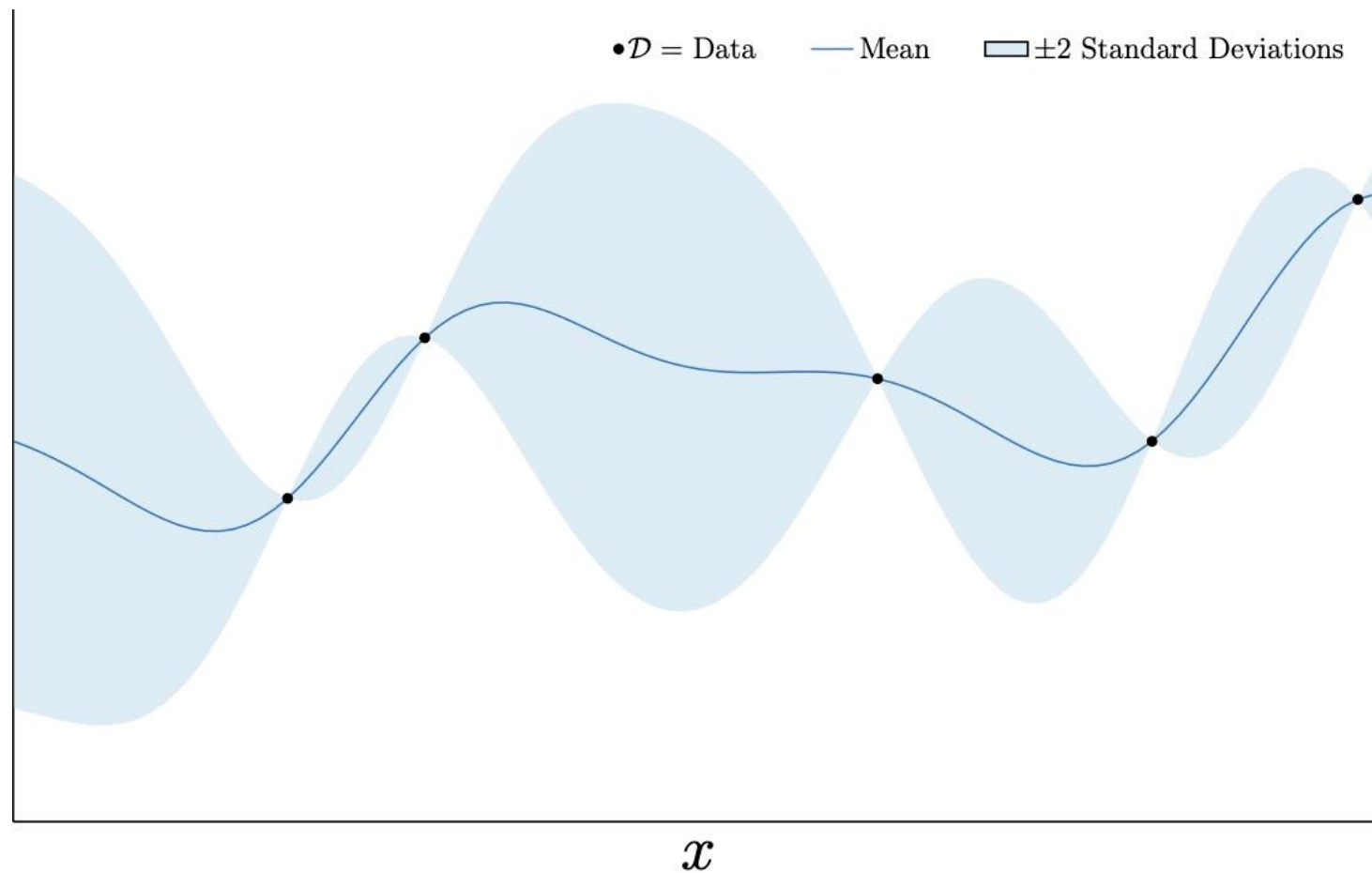
GP Prior

$$f: \mathbb{R}^p \mapsto \mathbb{R} \sim \mathcal{GP}(f; \mu(x) = 0, \Sigma(x, x') = \exp(-(x - x')^2))$$

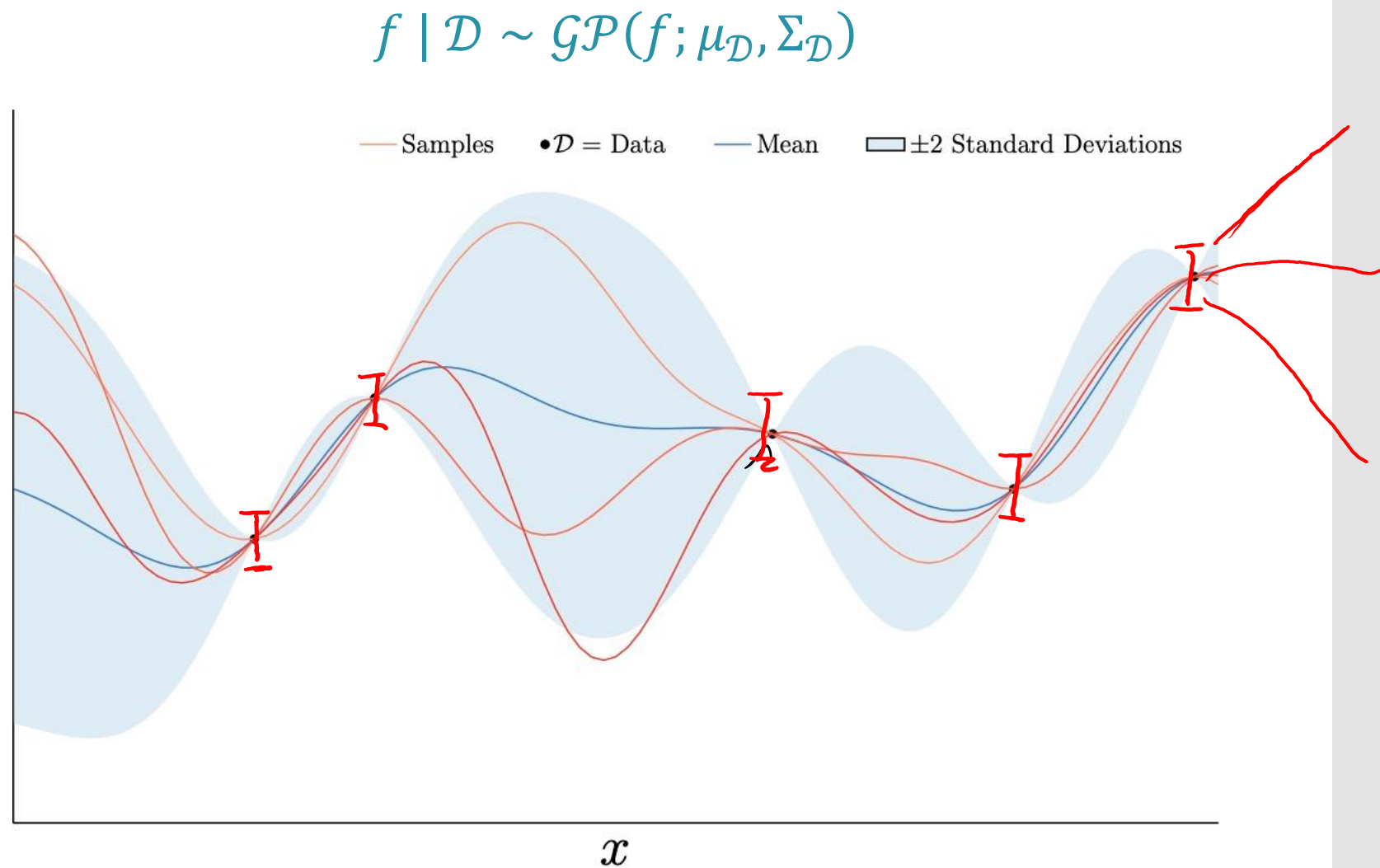


GP Posterior

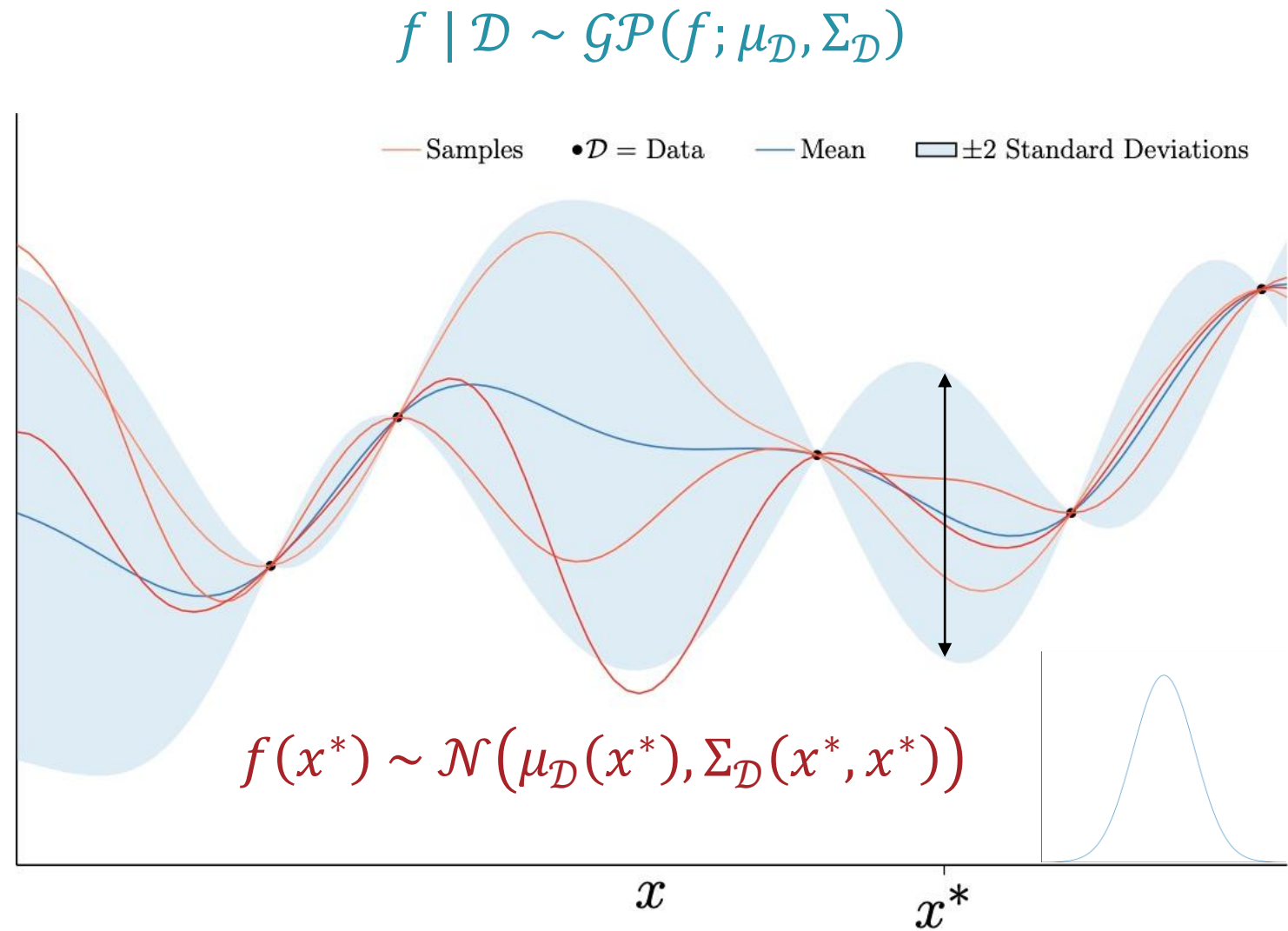
$$f \mid \mathcal{D} \sim \mathcal{GP}(f; \mu_{\mathcal{D}}, \Sigma_{\mathcal{D}})$$



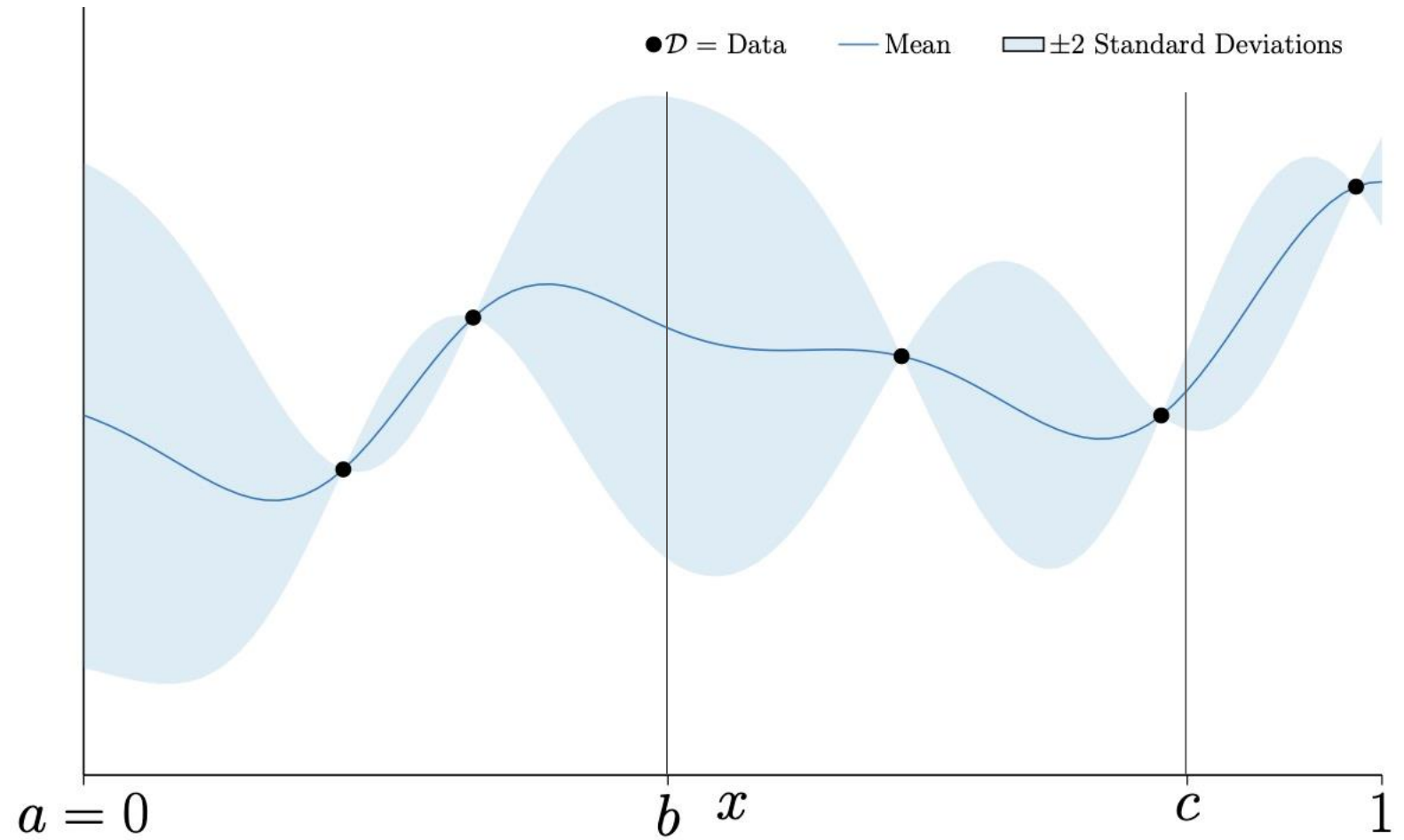
GP Posterior



GP Posterior

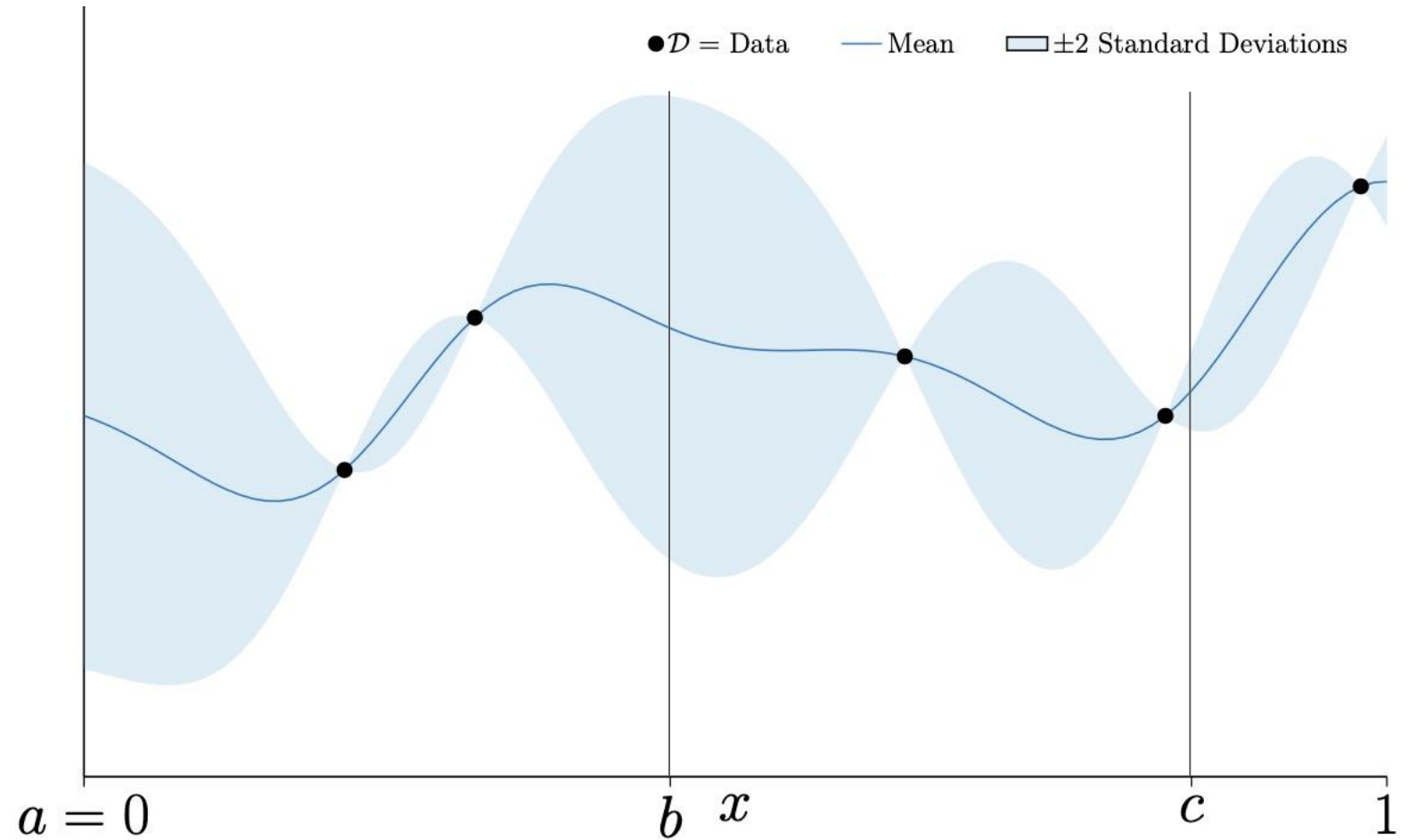


Active Learning



Suppose you can add one data point to your training data.

Which point would you add and why?

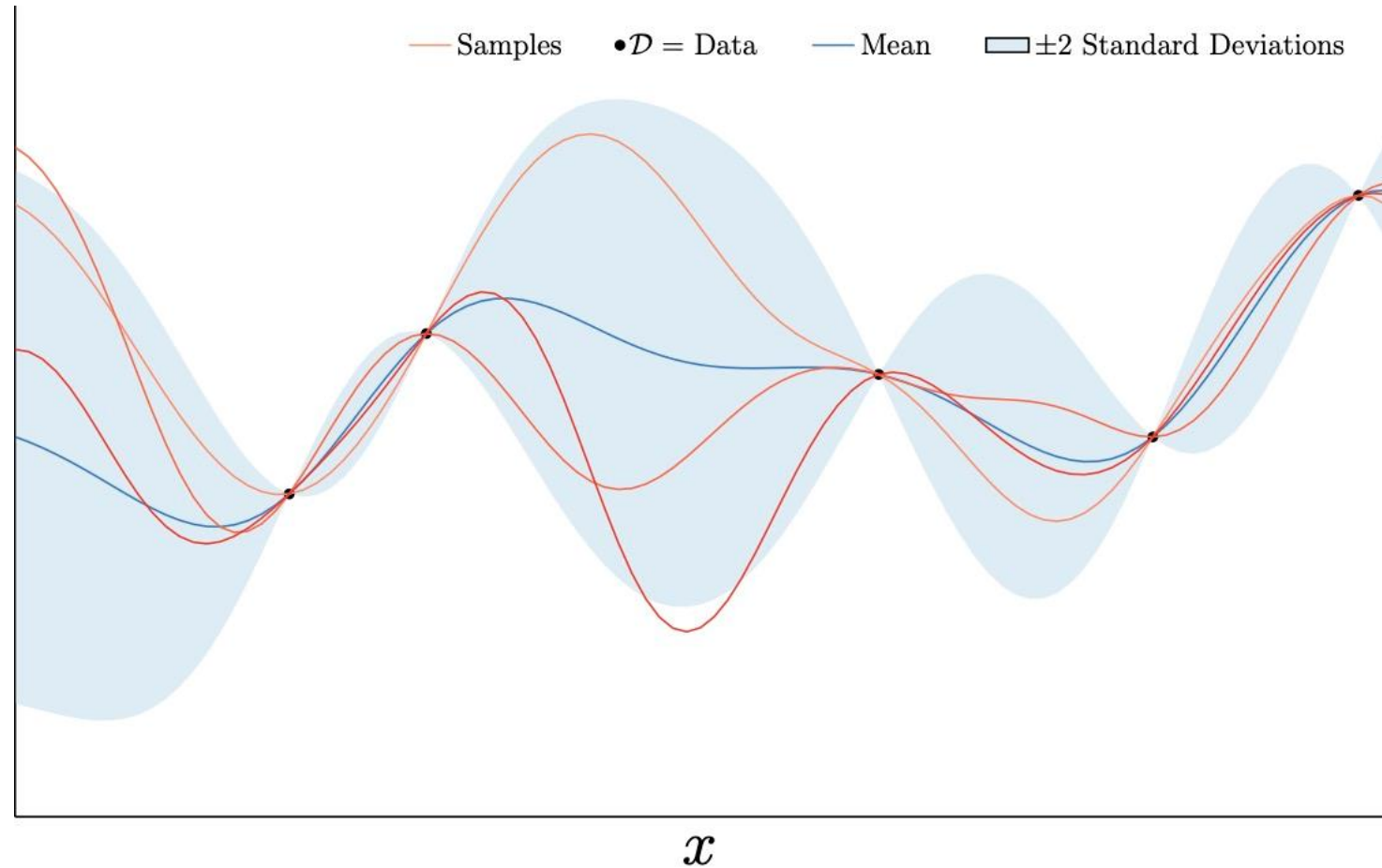


Are GPs:

1. parametric or nonparametric

2. generative or discriminative

$$f \mid \mathcal{D} \sim \mathcal{GP}(f; \mu_{\mathcal{D}}, \Sigma_{\mathcal{D}})$$

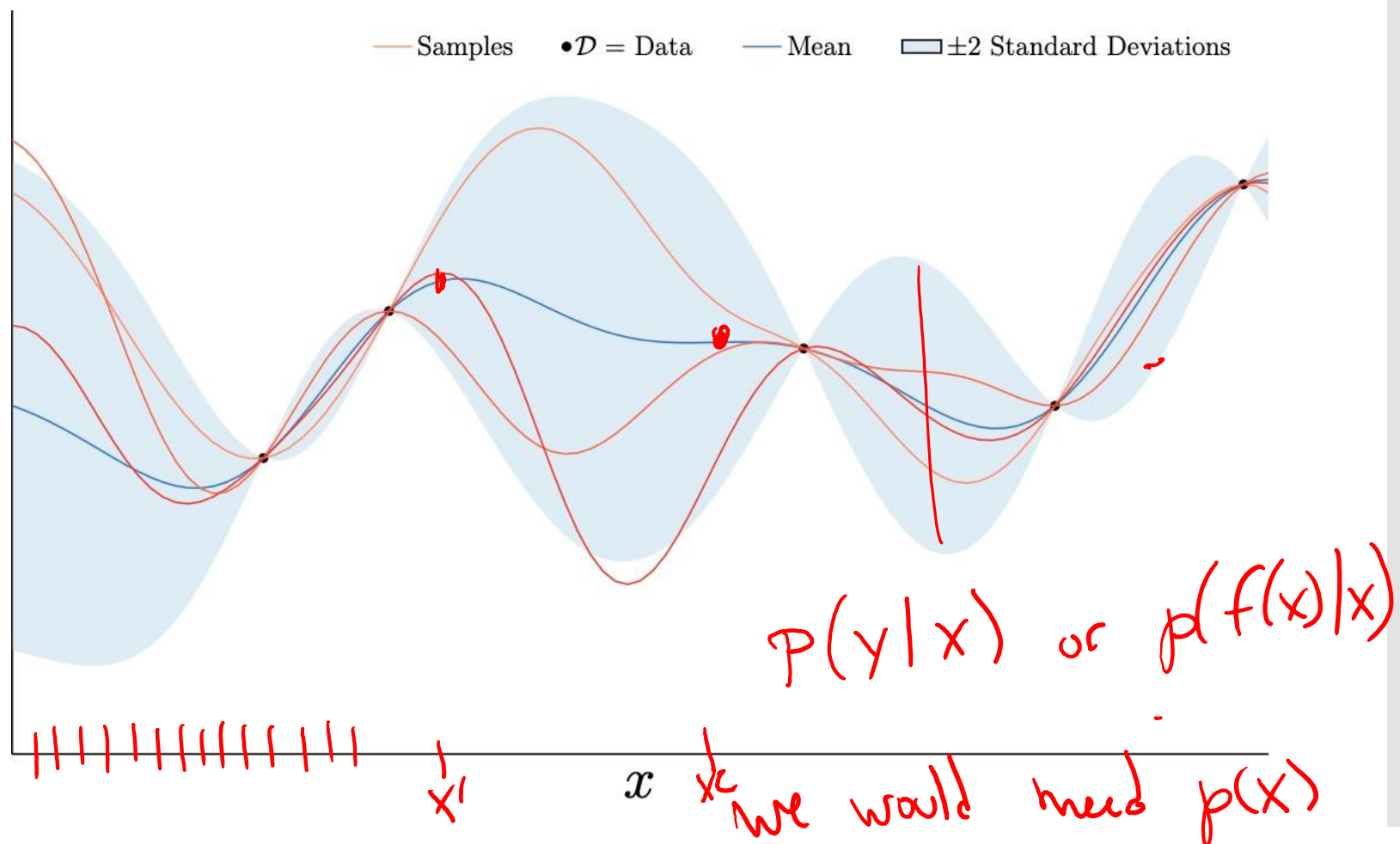


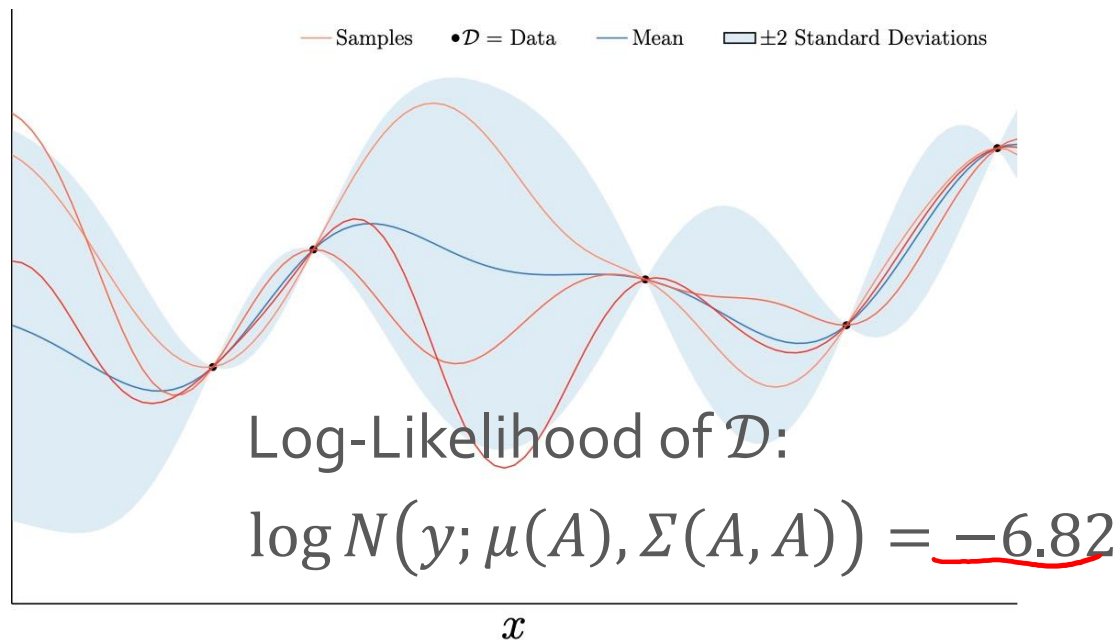
GPs are

1. parametric or nonparametric

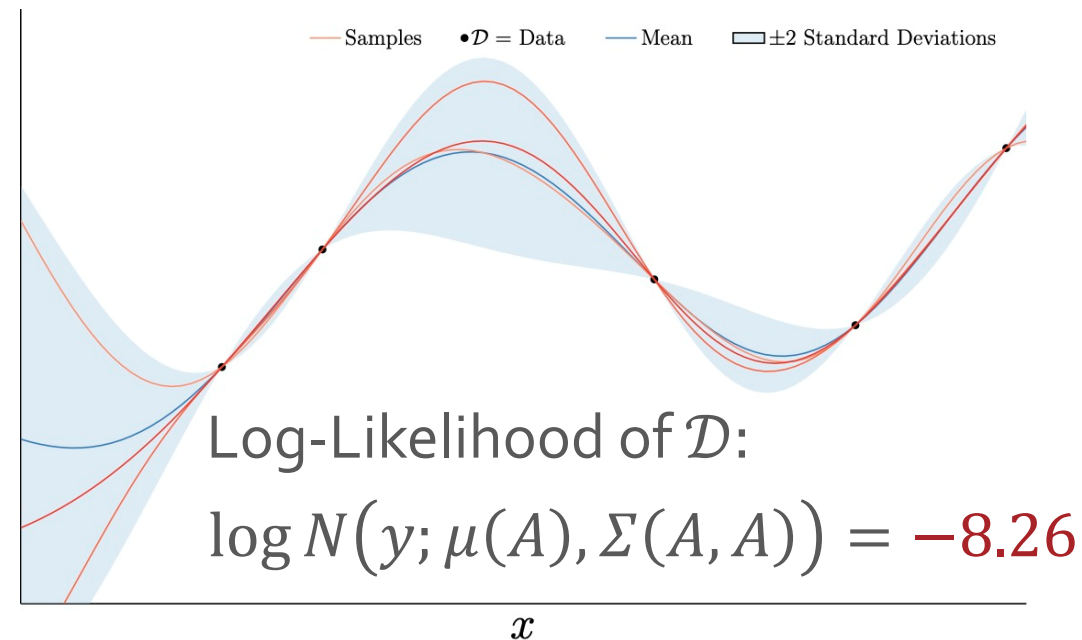
2. generative or discriminative

$$f \mid \mathcal{D} \sim \mathcal{GP}(f; \underline{\mu_{\mathcal{D}}, \Sigma_{\mathcal{D}}})$$





$$f \sim \mathcal{GP}\left(f; 0, \underline{(1^2)} \exp\left(-\frac{(x - x')^2}{\underline{1^2}}\right)\right)$$



$$f \sim \mathcal{GP}\left(f; 0, \underline{(2^2)} \exp\left(-\frac{(x - x')^2}{\underline{2^2}}\right)\right)$$

Kernel Hyperparameters