

Summary on Statistics

1) Useful probability dists

1.1. Binary variables

$$X = \{0, 1\}$$

$$\text{Let } P(X=1) = \mu$$

$$\hookrightarrow \text{Bern}(X|\mu) = \mu^x (1-\mu)^{1-x}$$

$\hookrightarrow$ Basis of binary classification

$$\text{Dataset } D = \{x_i\}_{i=1}^N \quad (\text{i.i.d.})$$

$$\downarrow$$
$$\text{Likelihood } P(D|\mu) = \prod_{i=1}^N \text{Bern}(x_i|\mu)$$

$\hookrightarrow$ object of interest

1.2. Categorical variables

$$X = \{1, 2, \dots, K\}$$

↙ parametrise it in the 1-to-K scheme
"one-hot encoding"

A point $x_i \rightarrow \vec{x}_i = \{x_i^1, x_i^2, \dots, x_i^K\}$

↙ binary vector

e.g. $K=5$

$$x_i = 3 \rightarrow \vec{x}_i = (0, 0, 1, 0, 0)$$

Distrib. is the Categorical distrib.

(or generalized Bernoulli)

$$g_{\text{Ben}}(\vec{x} | \vec{\mu}) = \prod_{k=1}^K \mu_k^{x_k}$$

↙
 $\vec{\mu} = \{\mu_k\}_{k=1}^K$

e.g. for
 $K=2$ classes

$$g_{\text{Ben}} = \mu_1^0 \mu_2^1 \quad \text{if} \quad \vec{x} = (0, 1)$$

$$= \mu^0 (1-\mu)^1 \Rightarrow \text{Bernoulli}$$

$$\mu_1 + \mu_2 = 1$$

$$\text{If } D = \{\vec{x}_i\}_{i=1}^N$$

PDF:

$$p(D|\vec{\mu}) = \prod_{i=1}^N \prod_{k=1}^K \mu_k^{x_{ik}}$$

$$-\ln p(D|\vec{\mu}) = -\sum_{i=1}^N \sum_{k=1}^K x_{ik} \ln \mu_k$$

↳ "cross-entropy" loss function

1.3) Continuous variables

Gaussian for a multivariate point $\vec{x}$ $\leftarrow$
dim = D

$$N(\vec{x}|\vec{\mu}, \Sigma) = \frac{1}{(2\pi)^{D/2}} \frac{1}{|\Sigma|^{1/2}} \times \exp \left\{ -\frac{1}{2} (\vec{x} - \vec{\mu})^T \Sigma^{-1} (\vec{x} - \vec{\mu}) \right\}$$

covariance

$\Sigma = D \times D$ matrix

$$\Sigma = \begin{pmatrix} \mathbb{E}[(x_1 - \mu_1)(x_1 - \mu_1)] & \dots \\ \mathbb{E}[(x_2 - \mu_2)(x_1 - \mu_1)] & \dots \\ \vdots & \\ \mathbb{E}[(x_D - \mu_D)(x_1 - \mu_1)] & \dots \end{pmatrix}$$

2.) Frequentist approach to inference

Likelihood of data D

$$p(D|\mu) = \prod_{i=1}^N p(x_i|\mu)$$

The aim is to model μ

$$\mu = h(\vec{w})$$

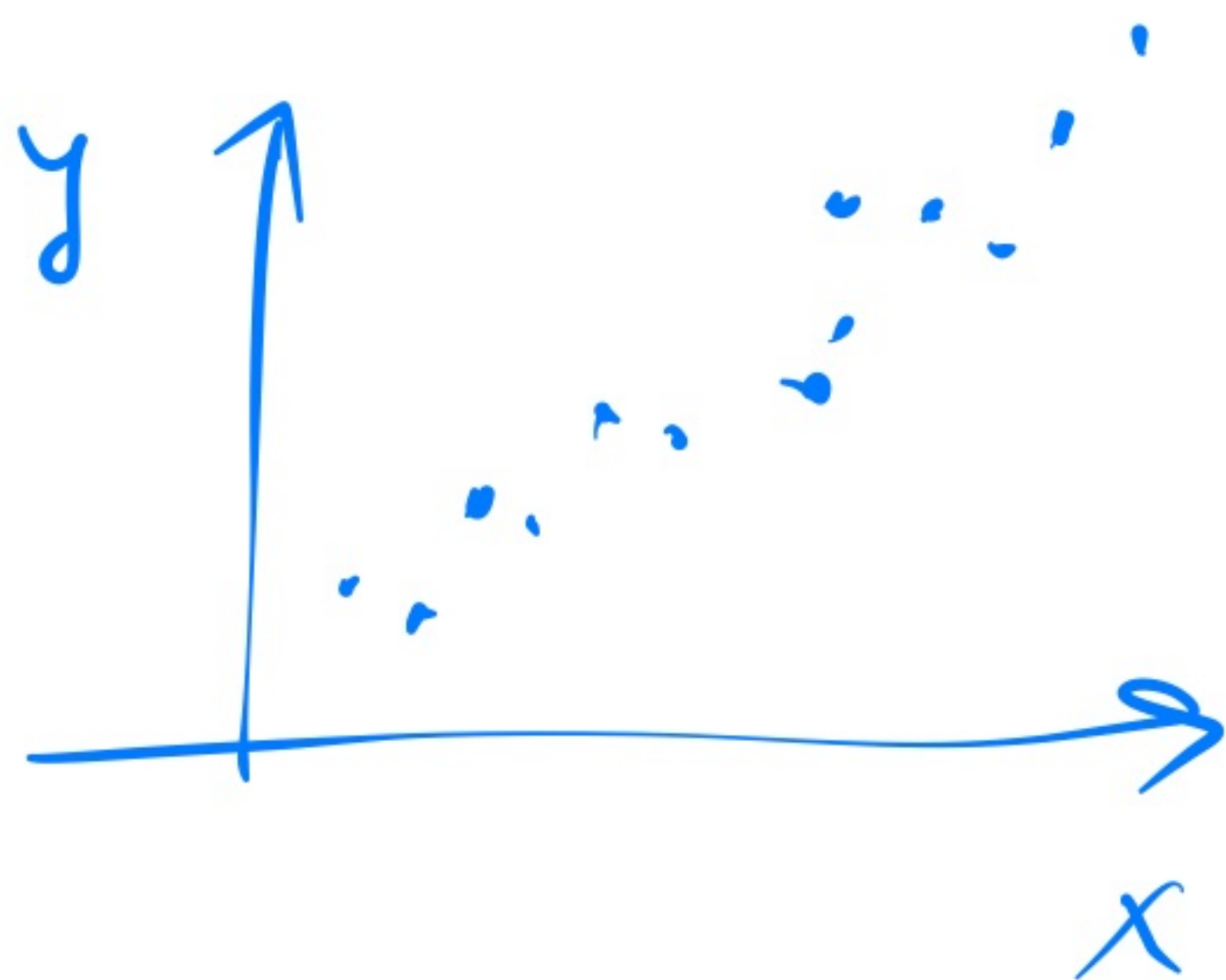
↳ predefined function

$$\vec{w}_{MLE} = \operatorname{argmax} p(D|\mu) = \mathcal{L}(\vec{w})$$

For example

$$D = \{x_i, y_i\}_{i=1}^N$$

↳ output



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

let's model

$$\mu(\vec{x}, \vec{w}) = ax + b$$

$$\vec{w} = \{a, b\} \quad \text{where } f(x, \vec{w})$$

Likelihood

$$\mathcal{L}(a, b) = \prod_{i=1}^N \mathcal{N}(y_i | f(x_i, a, b), \sigma^2)$$

$$-2 \ln \mathcal{L} = -2 \sum_i \left\{ \ln \frac{1}{\sqrt{2\pi}} + \frac{1}{2} \ln \frac{1}{\sigma^2} - \frac{(y_i - ax_i - b)^2}{2\sigma^2} \right\}$$

maximise $\mathcal{L}$

equivalent to minimise

$$\Rightarrow \mathcal{E} = \sum_i (y_i - ax_i - b)^2$$

obtain

least squares fitting procedure

$$a^* = \frac{N \sum_i y_i x_i - (\sum_i x_i)(\sum_i y_i)}{N \sum_i x_i^2 - (\sum_i x_i)^2}$$

$$b^* = \frac{1}{N} \left(\sum_i y_i - a^* \sum_i x_i \right)$$

Once we have $\vec{w}_{MLE}$

$\Rightarrow$ prediction for a new point

$$p(y_{\text{new}} | \vec{x}_{\text{new}}, \vec{w}_{\text{MLE}})$$

$$\uparrow = \mathcal{N}(y_{\text{new}} | a^* x_{\text{new}} + b^*, \sigma^2)$$

point-like prediction without associated uncertainties



a way out $\Rightarrow$ bootstrapping

↓ Bayesian approach

The final aim typically is the estimation of predictive distribution

$$p(y_{\text{new}} | x_{\text{new}}, D) \quad \xrightarrow{\{x, y\} \text{ training}}$$

$$= \int d\vec{w} \underbrace{p(y_{\text{new}} | \vec{x}_{\text{new}}, \vec{w})}_{\text{likelihood of } y_{\text{new}}} \underbrace{p(\vec{w} | D)}_{\text{posterior}}$$

one can compute the 2nd moment in order to estimate the uncertainties of the prediction

- Assume some data about observable y
 $\rightarrow$ take it as following a Gaussian distn

$$\Rightarrow P(\vec{y} | \vec{w}) = N(\vec{y} | X \cdot \vec{w} + \vec{b}, L^{-1})$$

$\downarrow$
 inverse of the
 cov. matrix
 = "precision matrix"

- Assume that $\vec{w}$ follow a Gaussian as well { prior on $\vec{w}$ }

$$\Rightarrow P(\vec{w}) = N(\vec{w} | \vec{\mu}, \Lambda^{-1})$$

Question: $P(\vec{w} | \vec{y})$? (posterior)

assignment $P(\vec{y})$? ("evidence")

Answer: $P(\vec{y}) = N(\vec{y} | X\vec{\mu} + \vec{b}, L^{-1} + X\Lambda^{-1}X^T)$

$$P(\vec{w} | \vec{y}) = N(\vec{w} | \Sigma (X^T L (\vec{y} - \vec{b}) + \Lambda \vec{\mu}), \Sigma)$$

Σ

$\uparrow$
 cov. of
 $\Sigma = (I + X^T L X)^{-1}$