

$$\begin{bmatrix} ML \end{bmatrix}$$

Supervised learning

$$D = \frac{1}{2} \sum_{n=1}^N (\underline{\tilde{x}}_n, y_n)$$

input output

$$p(y|x)$$

$$f(\bar{x}) \approx y$$

$$f(\overline{x}) \in [0, 1]$$

$$\hat{p}(\text{cat}(\bar{x}))$$

Unsupervised learning

$$D = \{ \bar{x}_n \}_{n=1}^N$$

dimensionality reduction

feature extraction

$$\bar{z} \in \mathbb{R}^d$$

represent. learning

1 Mpix $3 \cdot 10^6$
 $\bar{x} \in \mathbb{R}$

$\bar{x} \in \mathbb{R}$

500
R

Clustering

Wichtig

$\bar{z} \in \mathbb{R}^d$
 $d=0$

$$d=0$$

Autoencoder

Project

cat

422

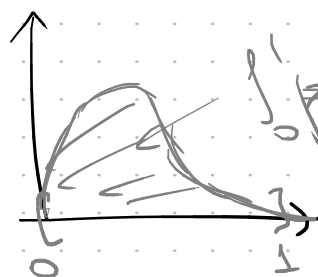
$$P(\overline{X})$$

$$X \sim p(\mathbb{Z})$$

random noise

$$L = \left\| \frac{\hat{x}}{\hat{\sigma}} - \bar{x} \right\|^2$$

$$\phi(x)$$

$$p(x)$$


$$\int_0^1 f(x) dx = 1$$

Joint prob

$P(x, y)$

$$(\neg(x=a) \wedge \neg(y=b))$$

$$P(x) = \sum_y P(x, y)$$

cond. prob

$$p(x|y)$$

$$p(x=a|y=b)$$

$$p(x|y) = \frac{p(x,y)}{p(y)}$$

$$p(x,y) \stackrel{\text{def}}{=} \underbrace{p(x|y)} \underbrace{p(y)}$$
$$\stackrel{\text{def}}{=} p(x)p(y|x)$$

$$p(y|x) = \frac{p(x|y)p(y)}{p(x)}$$

Bayes rule formula

independence

$$p(x,y) = p(x)p(y)$$

$$p(x,y|z) = p(x|z)p(y|z)$$

Conditional independence

$$p(x|y) = p(x) \rightarrow \text{max}$$

posterior distribution

$$p(\bar{\theta}|\underline{D}) = \frac{p(\underline{D}|\bar{\theta})p(\bar{\theta})}{p(\underline{D})}$$

data

model parameters

max $\bar{\theta}$

likelihood

$$p(\underline{D}|\bar{\theta}) \cdot p(\bar{\theta})$$

data

prior distribution

data

evidence prob.

max $\bar{\theta}$

$$p(\bar{\theta}|\underline{D}) \propto p(\underline{D}|\bar{\theta})p(\bar{\theta})$$

proportional

$$p(\underline{D}) = \int p(\bar{\theta})p(\underline{D}|\bar{\theta})d\bar{\theta}$$

$$\log p(\bar{\theta}|\underline{D}) = \log p(\underline{D}|\bar{\theta}) + \log p(\bar{\theta})$$

term

hthhht

$$p(\underline{D}|\bar{\theta}) =$$

$$= p(d_1|\bar{\theta})p(d_2|\bar{\theta}) \dots p(d_n|\bar{\theta})$$

Example: coin tossing

$$\theta = p(\text{heads}|\theta)$$

$$1-\theta = p(\text{tails}|\theta)$$

$$D = hthhltt$$

$$p(D|\theta) = p(h|\theta)p(t|\theta) \dots p(t|\theta) = \theta^4 (1-\theta)^3 \quad \text{likelihood}$$

1) Maximum Likelihood

$$\theta_{ML} = \arg \max_{\theta} p(D|\theta)$$

$D = \begin{matrix} n \text{ heads} \\ m \text{ tails} \end{matrix}$

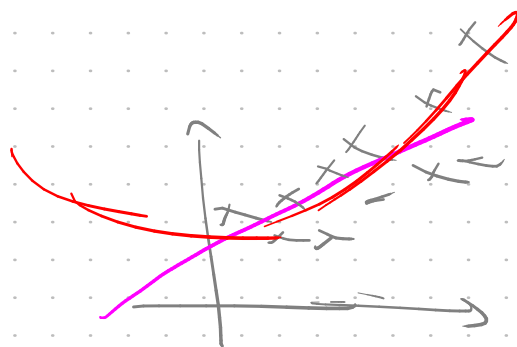
$$p(D|\theta) = \theta^n (1-\theta)^m \rightarrow \max ?$$

$$\frac{\partial p(D|\theta)}{\partial \theta} = n \theta^{n-1} (1-\theta)^m - m \theta^n (1-\theta)^{m-1} = 0$$

$$\theta^{n-1} (1-\theta)^{m-1} (n - (n+m)\theta) = 0$$

$$\theta \neq 0, \theta \neq 1, \quad \theta = \frac{n}{n+m} \quad n - (n+m)\theta$$

$$\theta_{ML} = \frac{n}{n+m} \rightarrow \text{true } \theta \quad n, m \rightarrow \infty$$

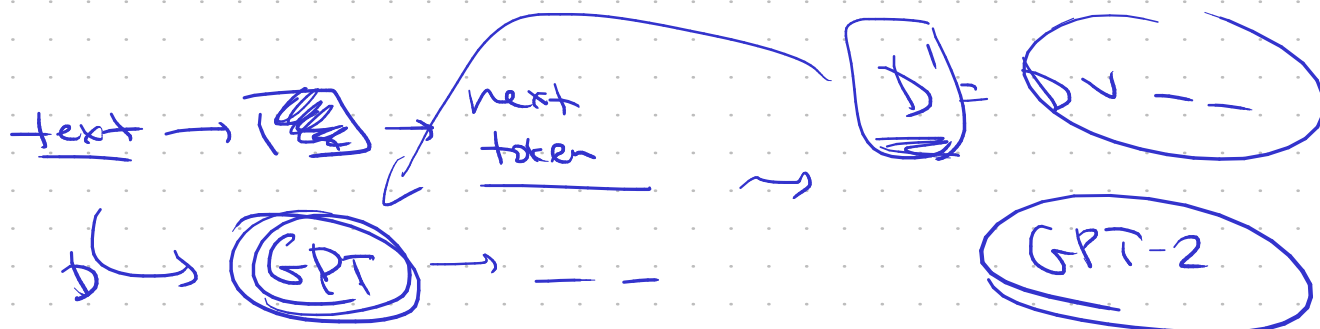


$$D = t \quad p(D|\theta) = 1 - \theta$$

$$\theta_{ML} = ?$$

$$\theta_{ML} = -\infty$$

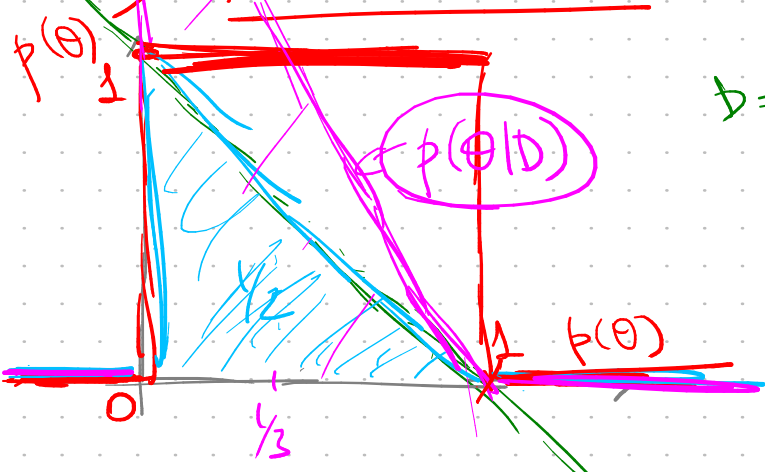
$$\theta_{ML} = 0$$



$$\theta_{max}$$

150M

2) Prior distribution

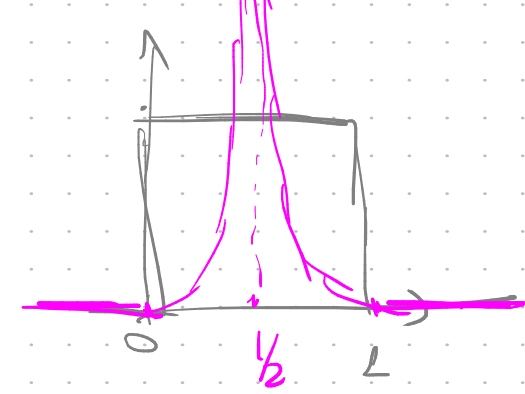


$$p(\theta)$$

$$D=t$$

$$p(D|\theta)=1-\theta$$

$$p(\theta|D) \propto p(\theta)p(D|\theta)$$



$$p(\theta|D) = \begin{cases} 2(1-\theta), & \theta \in [0,1] \\ 0, & \text{otherwise} \end{cases}$$

$$\theta_{MAP} = \arg \max_{\theta} p(\theta|D)$$

max a posteriori

$$\theta_{MAP} = \frac{n}{n+m}$$

$$p(\theta|D) \propto p(D|\theta)p(\theta) = \int_0^1 \theta^n (1-\theta)^m d\theta = \frac{1}{n+m+2}$$

3) Predictive distribution

$$E_{p(\theta|D)} [p(\text{heads}|\theta)]$$

$$p(\text{heads}|D) = \int p(\text{heads}|\theta) p(\theta|D) d\theta = \int_0^1 \theta \frac{\theta^n (1-\theta)^m}{\int_0^1 \theta^n (1-\theta)^m d\theta} d\theta = \frac{\int_0^1 \theta^{n+1} (1-\theta)^m d\theta}{\int_0^1 \theta^n (1-\theta)^m d\theta} = \frac{n+1}{n+m+2}$$

Laplace's rule of succession

$$p(\bar{\theta}|D) = \frac{p(D|\bar{\theta})p(\bar{\theta})}{p(D)}$$

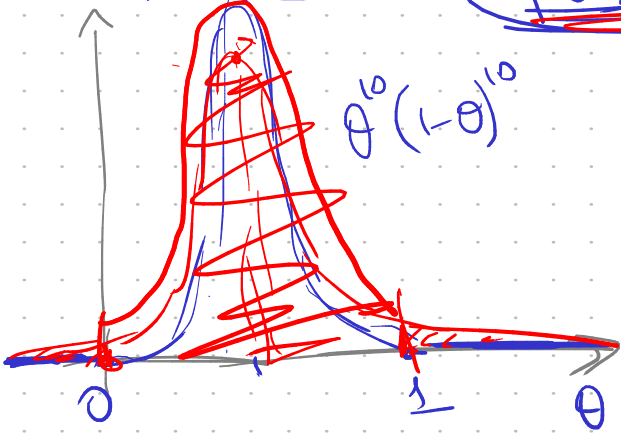
$p(\bar{x}|\bar{\theta})$ likelihood
posterior

pred. distr

$$p(\bar{x}|D) = \int p(\bar{x}, \bar{\theta}|D) d\bar{\theta} = \int p(\bar{x}|\bar{\theta}) p(\bar{\theta}|D) d\bar{\theta}$$

4) Prior?

$p(\theta) = ?$



$$p(\theta) = \begin{cases} c \cdot e^{-\frac{1}{2\sigma^2}(\theta - \frac{1}{2})^2} & [0, 1] \\ 0, & \text{---} \end{cases}$$

$p(\theta) \propto p(D|\theta) =$
 $= c \cdot e^{-\frac{1}{2\sigma^2}(\theta - \frac{1}{2})^2} \cdot \theta^n (1-\theta)^m$

likelihood

prior

posterior

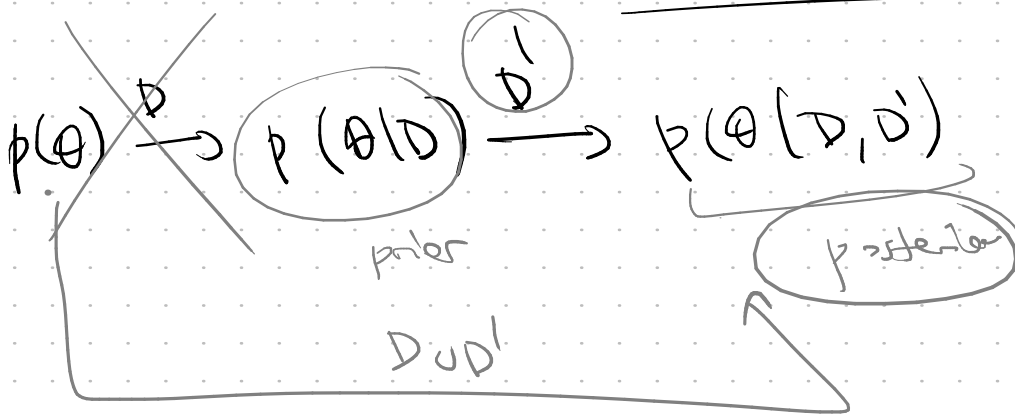
$\theta^n \cdot (1-\theta)^m$

$\times c \cdot \theta^{\alpha-1} (1-\theta)^{\beta-1}$

$\propto \theta^{(\alpha+n)-1} (1-\theta)^{(\beta+m)-1}$

conjugate prior

$D = (n, m) \times \text{Beta}(\theta | \alpha, \beta) \propto \text{Beta}(\theta | \alpha+n, \beta+m)$



$p(\theta|D) = \frac{p(D|\theta)p(\theta)}{p(D)}$

Laplace's rule

$p(\text{heads} | n, m) = \frac{n+1}{n+m+2}$

$p(D|\theta) \xrightarrow{\theta} \max$

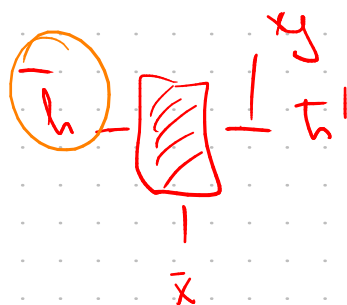
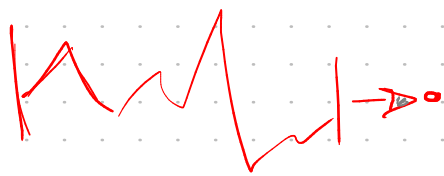
$\theta_{ML} = \frac{n}{n+m}$

$p(\theta|D) \xrightarrow{\theta} \max$

$\theta_{MAP} = \frac{n}{n+m}$

uniform $p(\theta)$

$p(x|D) = \int p(x|\bar{\theta})p(\bar{\theta}|D)d\bar{\theta} = E_{p(\bar{\theta}|D)}[p(x|\bar{\theta})]$



$$D = \{(\bar{x}_n, y_n)\}$$