

①

posterior distr. data

prior distr.

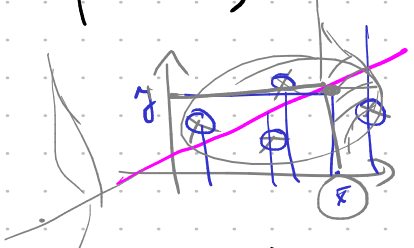
likelihood

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

params

(1) $p(D | \bar{\theta}) \xrightarrow{\bar{\theta}} \max, \bar{\theta}_{ML}$ - maximum likelihood

(2) $p(\bar{\theta} | D) \xrightarrow{\bar{\theta}} \max, \bar{\theta}_{MAP}$ - maximum a posteriori



(3) $p(\bar{x} | D)$ - predictive distribution

$$p(\bar{x} | D) = \int p(\bar{x}, \bar{\theta} | D) d\bar{\theta} = \int \underbrace{p(\bar{x} | \bar{\theta}, D)}_{\text{likelihood}} \underbrace{p(\bar{\theta} | D)}_{\text{posterior}} d\bar{\theta} =$$

$$p(D | \bar{\theta}) = \prod_{n=1}^N p(d_n | \bar{\theta})$$

$$= \underbrace{E_{p(\bar{\theta} | D)} [p(\bar{x} | \bar{\theta})]}_{\text{predictive distribution}}$$

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

② Bernoulli trials

$D = \text{hhtth}$, $\theta = p(\text{heads})$

$p(D | \theta) = \theta^m (1 - \theta)^n$, $m = \# \text{ heads}$, $n = \# \text{ tails}$

$\xrightarrow{\theta} \max$

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

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

$$m - (m + n)\theta = 0$$

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

$\theta = 0, 1, \frac{m}{m + n}$

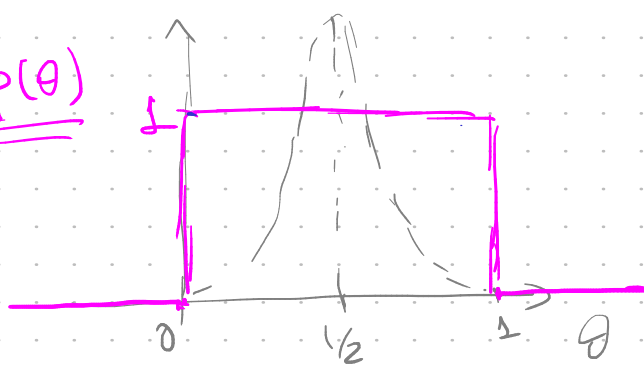
$D = t, \theta_{ML} = 0$

$\int_0^1 (1 - \theta) p(\theta | x) d\theta$

$p(D | \theta) = 1 - \theta \xrightarrow{\theta} \max, \theta \in [0, 1]$

$D = \text{hhh}, \theta_{ML} = 1$

p(θ)



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

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

$$\frac{p(y|x, D)}{1(w_0, w_2)}$$

$$p(\theta|D) = \begin{cases} \frac{1}{Z} \theta^m (1-\theta)^n, & \theta \in [0, 1] \\ 0, & \theta \notin [0, 1] \end{cases}$$

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

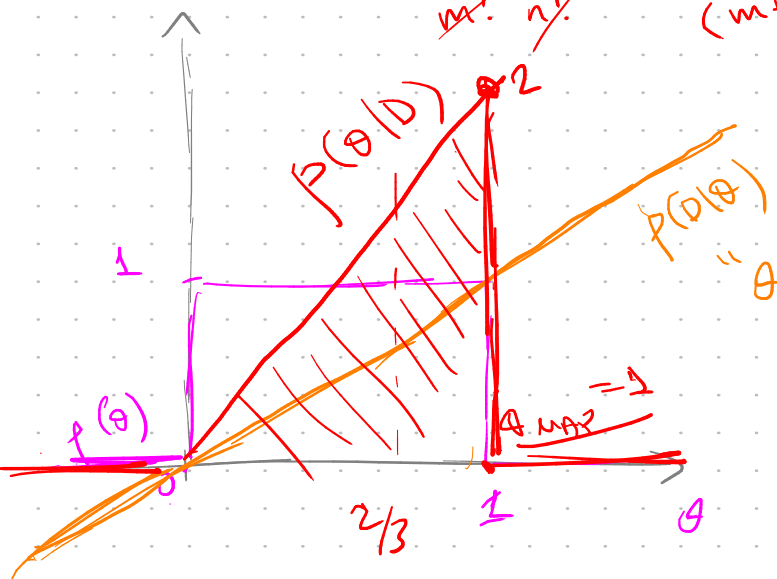
$$Z = \int_0^1 \theta^m (1-\theta)^n d\theta = B(m+1, n+1) = \frac{\Gamma(m+1) \Gamma(n+1)}{\Gamma(m+n+2)} = \frac{m! n!}{(m+n+1)!}$$

$$p(\theta_{MAP}|D) = \int p(\theta_{MAP}|\theta) p(\theta|D) d\theta =$$

$$= \int_0^1 \theta \cdot \frac{(m+n+1)!}{m! n!} \theta^m (1-\theta)^n d\theta = \frac{(m+n+1)!}{m! n!} \int_0^1 \theta^{m+1} (1-\theta)^n d\theta =$$

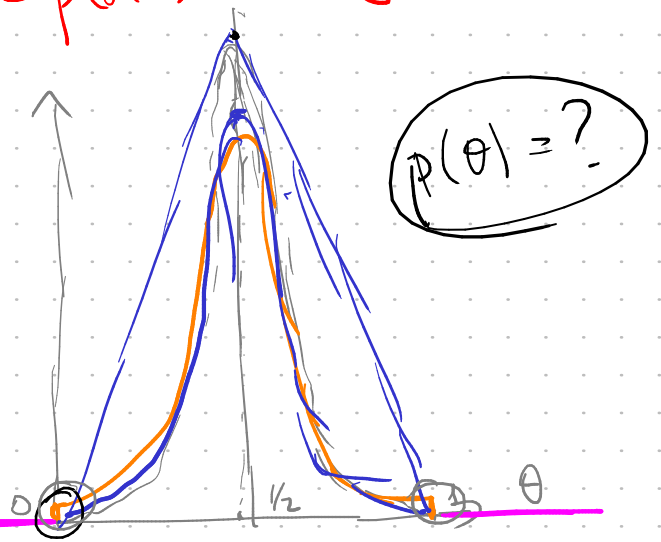
$$= \frac{(m+n+1)!}{m! n!} \cdot \frac{(m+1)! n!}{(m+n+2)!} =$$

Laplace's Rule of Succession



Bayesian Smoothing

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



3

Аппроксимация гаусс-регенерация

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

$$\underbrace{p(\theta)}_{e^{-\frac{1}{2}(\theta - \frac{1}{2})^2}} \times p(D|\theta) \propto \underbrace{p(\theta|D)}_{\theta^m (1-\theta)^n e^{-\frac{1}{2}(\theta - \frac{1}{2})^2}}$$

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

$$\underline{p(\theta) = \text{Beta}(\theta | \alpha, \beta) = \frac{1}{B(\alpha, \beta)} \cdot \theta^{\alpha-1} (1-\theta)^{\beta-1}}$$

Conjugate priors - convenient prior. $p(\bar{\theta} | \bar{x})$, τ, γ, ρ .
 $\uparrow$
 hyperparameter
 $p(\bar{\theta} | \bar{x}) \cdot p(D|\bar{\theta}) \propto p(\bar{\theta} | \bar{x}')$

$$\text{Beta}(\theta | \alpha, \beta) \times p(\underbrace{m}_{n \times \text{heads}} \times \underbrace{n}_{n \times \text{tails}} | \theta) \propto \text{Beta}(\theta | \underbrace{\alpha+m}_{\beta+n})$$

$$p(\bar{\theta} | \bar{x}_0) \times p(D|\bar{\theta}) \propto p(\bar{\theta} | \bar{x}_1) \times p(D'|\bar{\theta}) \propto p(\bar{\theta} | \bar{x}_2)$$

$$p(D|\bar{\theta}) = \prod_{n=1}^N p(d_n | \bar{\theta}) \quad \theta^{10} (1-\theta)^{10}$$

$$\log p(D|\bar{\theta}) = \sum_{n=1}^N \log p(d_n | \bar{\theta})$$

$$\boxed{p(\theta) = \text{Beta}(\theta | \alpha, \beta)} \Rightarrow \text{posterior w/ unif. prior} \\
+ D = (\alpha-1, \beta-1) \\
\underbrace{\hspace{10em}}_{\text{equivalent sample size}} \\
\text{Unif}(0,1) \times \theta^{\alpha-1} (1-\theta)^{\beta-1} \\
p(\theta)''$$

$$p(\theta) = \text{Beta}(\theta | \frac{1}{2}, \frac{1}{2}) = \theta^{-\frac{1}{2}} (1-\theta)^{-\frac{1}{2}} = \frac{1}{\sqrt{\theta(1-\theta)}}$$

④ Multinomial distribution

$$p(D|\bar{\theta}) = p(n_1, n_2, \dots, n_k | \bar{\theta}) = \theta_1^{n_1} \theta_2^{n_2} \dots \theta_k^{n_k}$$

$$p(\bar{\theta}) = \text{Dir}(\bar{\theta} | \bar{\alpha}) = \frac{1}{\text{Dir}(\bar{\alpha})} \theta_1^{\alpha_1-1} \theta_2^{\alpha_2-1} \dots \theta_k^{\alpha_k-1}$$

$$\approx 1 - \theta_1 - \theta_2 - \dots - \theta_{k-1}$$

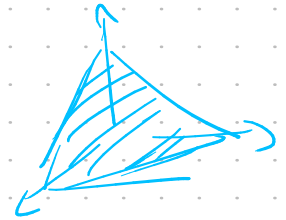
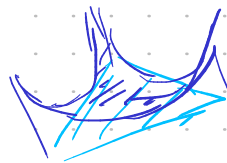
$$p(\bar{\theta}) = \text{Dir}(\bar{\theta} | \bar{\alpha} = (\frac{1}{10}, \frac{1}{10}, \dots, \frac{1}{10}))$$

Sparsity

LDA

$$\bar{\alpha} = \begin{pmatrix} \frac{1}{10} \\ \frac{1}{10} \\ \vdots \\ \frac{1}{10} \end{pmatrix}$$

$$\text{Dir}(\bar{\theta} | \bar{\alpha})$$



$\bar{d} \bar{t}$

$$p(d|t) = \frac{p(d)p(t|d)}{p(d)p(t|d) + p(\bar{d})p(t|\bar{d})} \approx \frac{1}{6}$$

$\approx 1/100$ $\approx 1/100$ $\approx 1/100$

$p(t)$