

Lecture 4: 1/30/17

Natural Data

naturally occurring
stochastic processes

Ex1: Coin Flipping
HTHTTTHH...

Ex2: Dice Rolls
3 2 2 1 6 5 4 1 2...

Ex3: Background Test taken by 601 Students



73, 100, 40, 68, 74, ...

Bernoulli($\phi = 0.5$)

Categorical($\vec{\theta}$)

$\theta_1 = 1/6, \theta_2 = 1/6 \dots \theta_6 = 1/6$

Gaussian(μ, σ^2)

$\mu = 70 \quad \sigma = 10$

Ex2: Dice Rolls

$\Pr(\text{die} = 1) = 1/6, \Pr(\text{die} = 2) = 1/6, \dots \Pr(\text{die} = 6) = 1/6$

$\Pr(\text{die} = 0) = 0, \Pr(\text{die} = 7) = 0$



event

probability

Generate Synthetic Data

Why?

- ① Plots for ML
- ② Helpful input when debugging ML
- ③ Helpful output of (certain) ML algos
- ④ Inference (see Gibbs Sampling)
- ⑤ Tells about the dist.

Def: Random variable X : value is the outcome of a random experiment.

$\Pr(X = x)$ — r.v.s by capital letters
— values by lowercase

event

$D = \{x^{(1)}, x^{(2)}, \dots, x^{(N)}\}$

If a r.v. X follows a Bernoulli w/parameter ϕ :

$X \sim \text{Bernoulli}(\phi)$

"is sampled from"

Categorical:

$$X \sim \text{Categorical}(\vec{\Theta}) \quad K = |\vec{\Theta}| \quad \vec{\Theta}^T = [\theta_1, \theta_2, \dots, \theta_K]$$

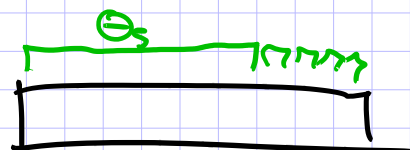
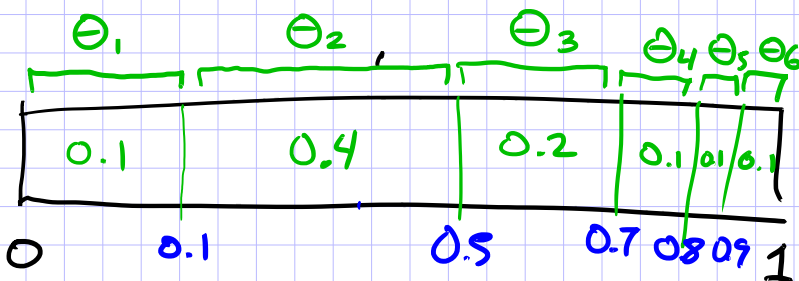
$$\sum_{k=1}^K \theta_k = 1, \quad 0 \leq \theta_k \leq 1, \quad \forall k$$

$$P_r(X=x) = \theta_x$$

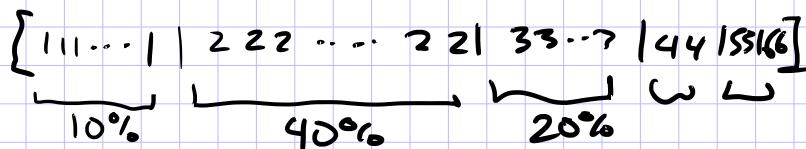
The ~~probability mass function~~ (pmf) for Cat:

$$p(x) \triangleq P_r(X=x) = \theta_x$$

Mult($\vec{\Theta}, N$)
where $N=1$
 $\equiv$
Cat($\vec{\Theta}$)



rand() $\leftarrow$ samples uniformly from $[0, 1]$



$$\text{rand}() = 0.4 \Rightarrow x^{(1)} = 2$$

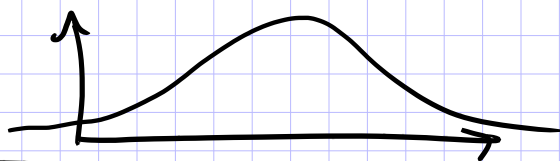
$$\text{rand}() = 0.6 \Rightarrow x^{(2)} = 3$$

$$\text{rand}() = 0.05 \Rightarrow x^{(3)} = 1$$

$\vdots$

Bernoulli: Just Categorical $K=2$ (coin flipping)

Gaussian:



$$X \sim \text{Gaussian}(\mu, \sigma^2)$$

$\uparrow$ continuous r.v.

Data Likelihood

Assume: Sample in D are iid

$\uparrow$ "independent and identically distributed"

① Indep. $P_r(X^{(1)}=6, X^{(2)}=3) = P_r(X^{(1)}=6) P_r(X^{(2)}=3)$

shortcut $\rightarrow P_r(X^{(1)}, X^{(2)}) = P_r(X^{(1)}) P_r(X^{(2)})$

$$\forall a, b \quad P_r(X^{(1)}=a, X^{(2)}=b) = P_r(X^{(1)}=a) P_r(X^{(2)}=b)$$

② Ident. Dist. each r.v. follows the same distribution

Ex: $X^{(i)} \sim \text{Categorical}(\vec{\Theta})$

$$P_r(X^{(i)}=x^{(i)}) = \theta_{x^{(i)}}, \quad \forall i$$

Ex: Dice Rolls

Q: What is the data likelihood of D under $\text{Cat}(\vec{\theta})$?

$$\begin{aligned} A: p(D|\vec{\theta}) &= p(X^{(1)}=x^{(1)}, X^{(2)}=x^{(2)}, \dots, X^{(N)}=x^{(N)}|\vec{\theta}) \\ &= p(X^{(1)}=x^{(1)}|\vec{\theta}) p(X^{(2)}=x^{(2)}|\vec{\theta}) \dots p(X^{(N)}=x^{(N)}|\vec{\theta}) \quad \text{by indep.} \\ &= \prod_{i=1}^N p(X^{(i)}=x^{(i)}) \\ &= \prod_{i=1}^N \theta_{x^{(i)}} \quad \text{by indep. dist.} \end{aligned}$$

Learning from Data

MLE: The principle of Maximum likelihood estimator (MLE):

Choose parameters that make the data "most likely".

Assumptions: Data generated iid from distribution $p^*(x|\vec{\theta}^*)$ and comes from a family of distr parameterized $\Theta \in \mathcal{H}$
 $\leftarrow$ set of possible parameters

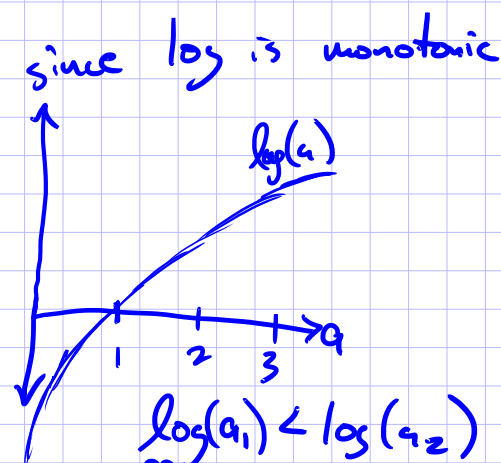
Formally:

$$\begin{aligned} \Theta_{MLE} &= \underset{\Theta \in \mathcal{H}}{\operatorname{argmax}} p(D|\Theta) \\ &= \underset{\Theta \in \mathcal{H}}{\operatorname{argmax}} \log p(D|\Theta) \\ &= \underset{\Theta \in \mathcal{H}}{\operatorname{argmax}} l(\Theta) \end{aligned}$$

usually a continuous optimization $\rightarrow$

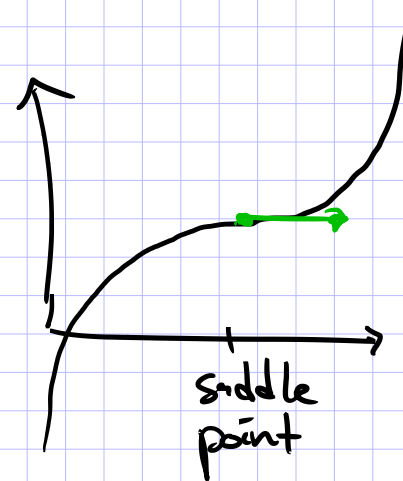
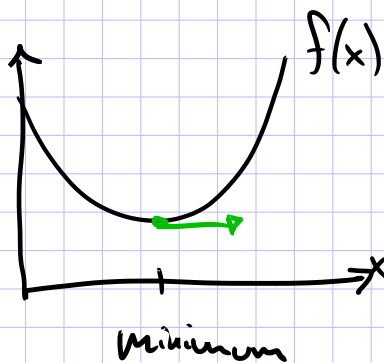
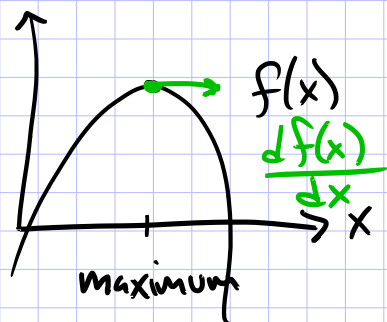
where $l(\Theta) \triangleq \log p(D|\Theta)$
 $\leftarrow$ "log-likelihood"

$\uparrow$ treat as function of Θ where D is constant



$$\begin{aligned} \log(a_1) &< \log(a_2) \\ \text{iff } a_1 &< a_2 \\ \Rightarrow \log(f(a_1)) &< \log(f(a_2)) \\ \text{iff } f(a_1) &< f(a_2) \end{aligned}$$

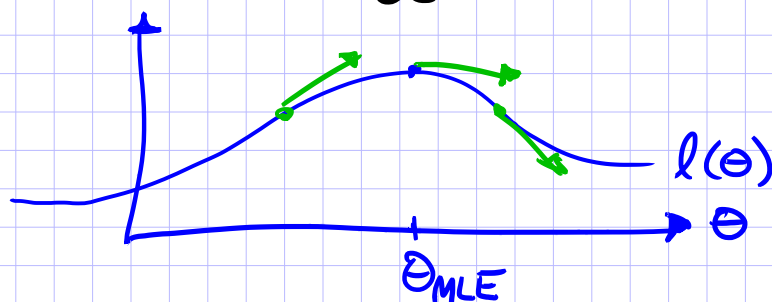
Zero Derivatives



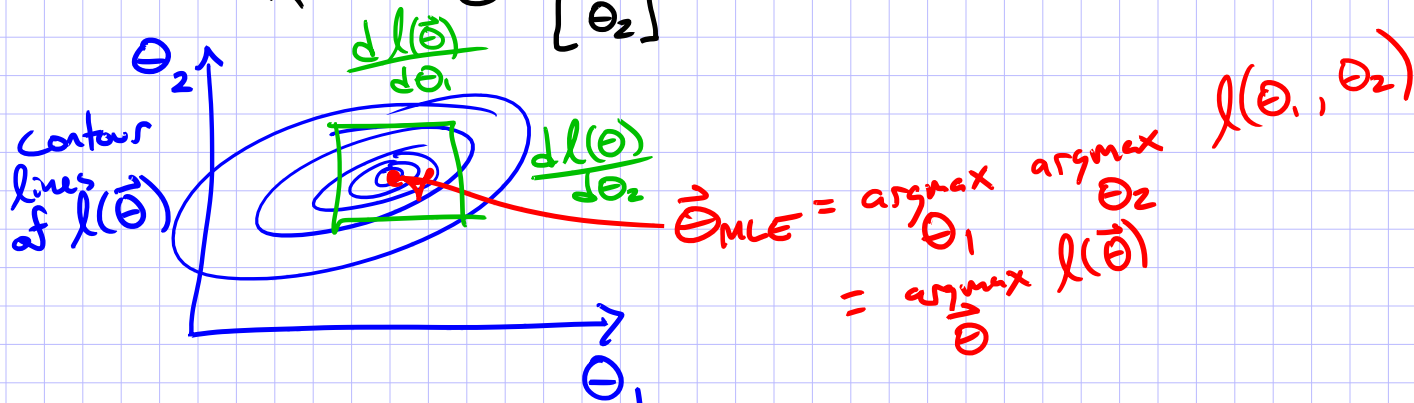
Optimization for MLE

Ex#1: Suppose $\Theta = \mathbb{R} \Rightarrow$ infinitely large

$$\Theta_{MLE} = \underset{\theta \in \Theta}{\operatorname{argmax}} l(\theta)$$



Ex#2: $\Theta = \mathbb{R}^2 \Rightarrow \vec{\theta} = \begin{bmatrix} \theta_1 \\ \theta_2 \end{bmatrix}$



$$\vec{\theta}_{MLE} = \underset{\theta_1}{\operatorname{argmax}} \underset{\theta_2}{\operatorname{argmax}} l(\theta_1, \theta_2) = \underset{\vec{\theta}}{\operatorname{argmax}} l(\vec{\theta})$$

Key idea: the top of any hill is flat!

Formally: $\frac{dl(\vec{\theta})}{d\theta_k} = 0 \quad \forall k$

Ex#3 MLE of Categorical

① Write $l(\theta)$

Let $N_j = \# \text{ times } j \text{ appears in } D$

Likelihood: $p(D|\vec{\theta}) = \prod_{i=1}^N \theta_{x^{(i)}} = \prod_{k=1}^K \theta_k^{N_k}$ exponent!!!

Log-likelihood:

$$l(\theta) = \log \prod_{k=1}^K \theta_k^{N_k} = \sum_{k=1}^K N_k \log \theta_k$$

② Set deriv to 0 $\downarrow$ $\frac{d l(\theta)}{d \theta_j} = \frac{d}{d \theta_j} \left(\sum_{k=1}^K N_k \log \theta_k \right) = \frac{d}{d \theta_j} N_j \log \theta_j = \frac{N_j}{\theta_j}$

$\leftarrow N_1 \log \theta_1 + N_2 \log \theta_2 + \dots + N_k \log \theta_k$

$$\frac{d l(\theta)}{d \theta_j} = 0 \Rightarrow \frac{N_j}{\theta_j} = 0 \Rightarrow \theta_j = +\infty \leftarrow \text{Wrong!!!}$$

WRONG

We forgot that $\sum_{k=1}^K \theta_k = 1$

The set of valid parameters

$$\mathcal{H} = \left\{ \vec{\theta} : \vec{\theta} \in \mathbb{R}^K \text{ and } \sum_{k=1}^K \theta_k = 1 \right\}$$

$$\theta_{MLE} = \underset{\theta \in \mathcal{H}}{\operatorname{argmax}} l(\theta)$$