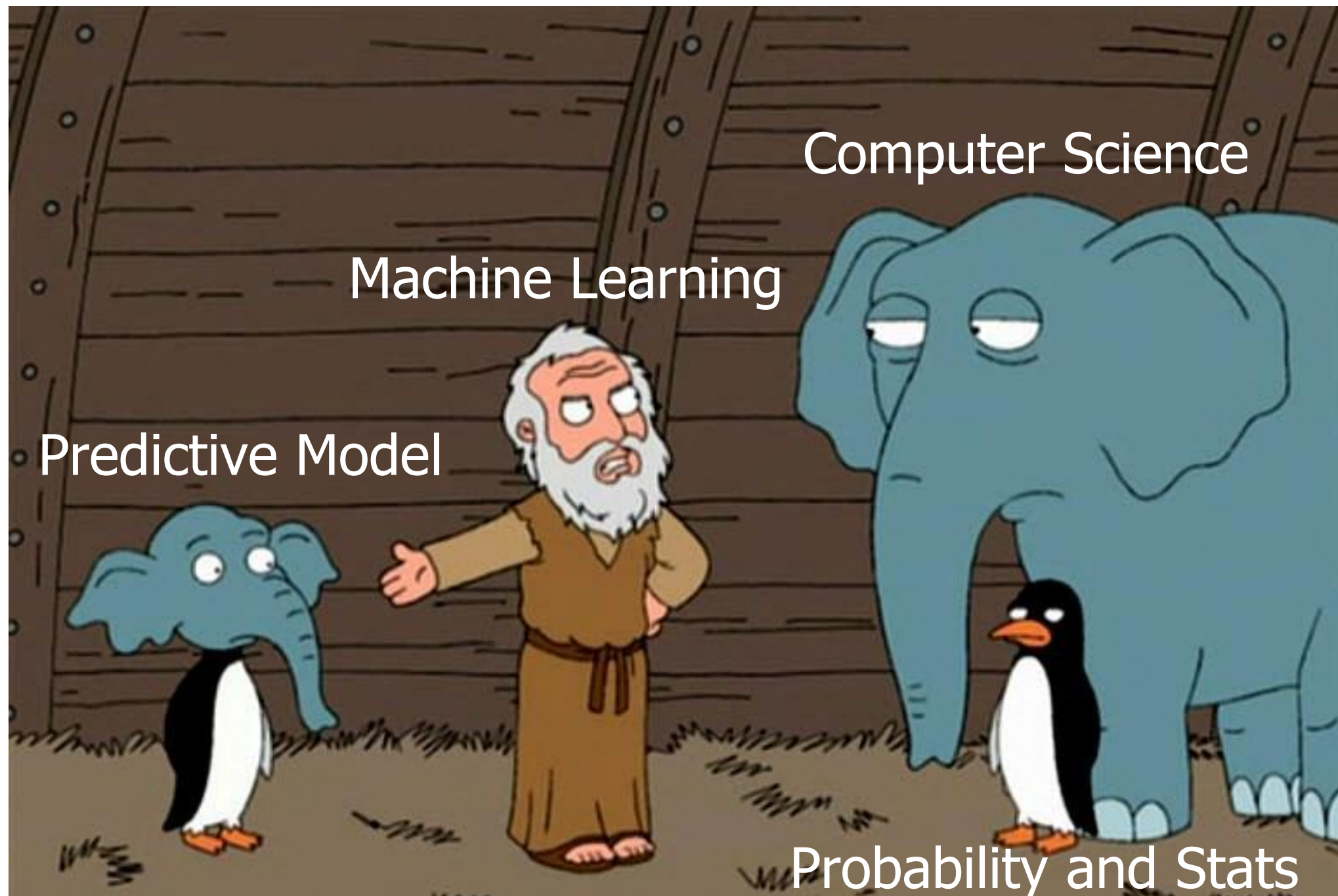




Computer Science

Machine Learning

Predictive Model

Probability and Stats

Probability and Statistics

Mahdi Roozbahani
Georgia Tech

Outline

- Probability Distributions 
- Joint and Conditional Probability Distributions
- Bayes' Rule
- Mean and Variance
- Properties of Gaussian Distribution
- Maximum Likelihood Estimation

Probability

- A **sample space S** is the set of all possible outcomes of a conceptual or physical, repeatable experiment. (S can be finite or infinite.)
 - E.g., S may be the set of all possible outcomes of a dice roll: S
(1 2 3 4 5 6) $P(X=x) = P(x=1) = \frac{1}{6}$
 - E.g., S may be the set of all possible nucleotides of a DNA site: S
(A C G T)
- E.g., S may be the set of all possible time-space positions of an aircraft on a radar screen.
- An **Event A** is any subset of S
 - Seeing "1" or "6" in a dice roll; observing a "G" at a site; UA007 in space-time interval



Three Key Ingredients in Probability Theory

A **sample space** is a collection of all possible **outcomes**

Random variables X represents **outcomes** in sample space

Probability of a random variable to happen $p(x) = p(X = x)$

$$p(x) \geq 0 \quad 0 \leq p(x) \leq 1$$

Pdf $\rightarrow$ probability distribution function

Pdf $\rightarrow$ Pmf

Continuous variable

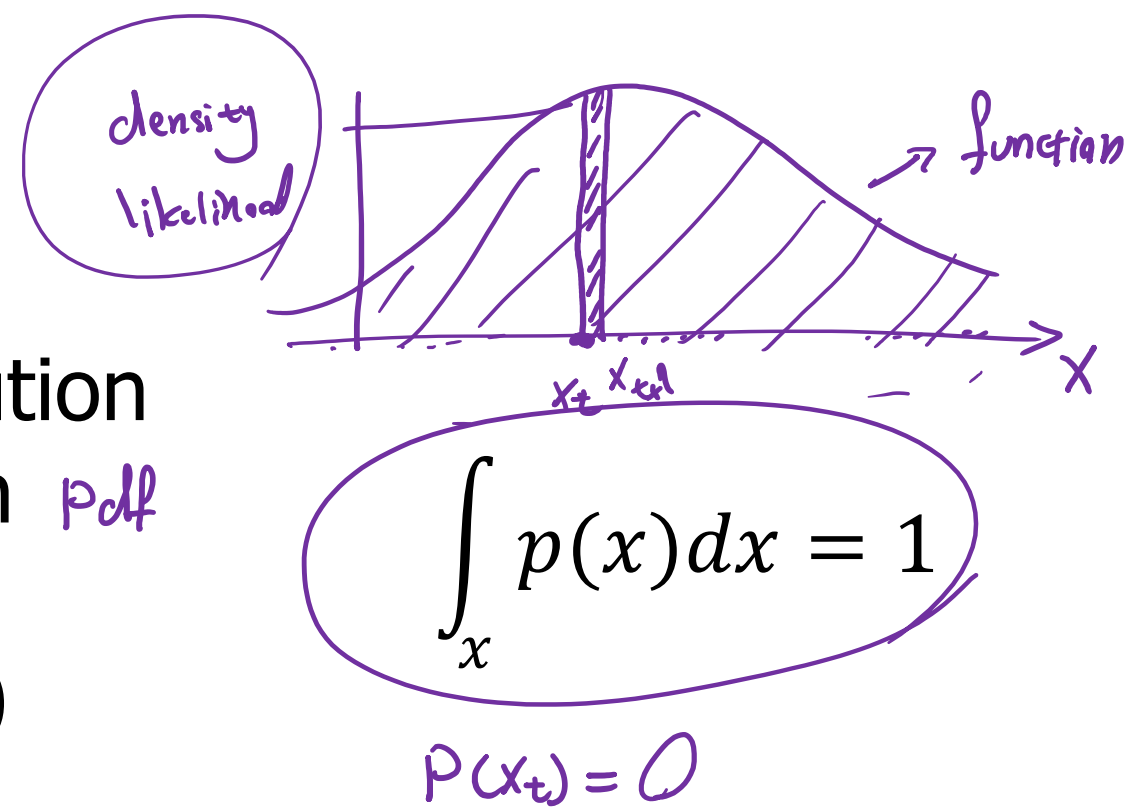
Continuous probability distribution

Probability density function pdf

Density or likelihood value

Temperature (real number)

Gaussian Distribution



Discrete variable

Discrete probability distribution

Probability mass function pmf

Probability value

Coin flip (integer)

Bernoulli distribution

probability

H T

$\frac{1}{2}$ $\frac{1}{2}$

$\frac{1}{2} + \frac{1}{2} = 1$

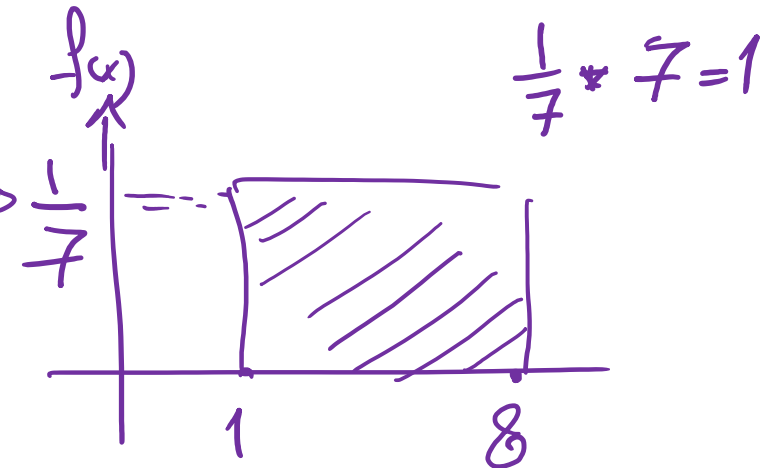
$$\sum_{x \in A} p(x) = 1$$

Continuous Probability Functions

- Examples: $\Theta \rightarrow \text{Parameters}$
 $\{b, a\} \in \Theta$

- Uniform Density Function:

$$f_x(x) = \begin{cases} \frac{1}{b-a} & \text{for } a \leq x \leq b \\ 0 & \text{otherwise} \end{cases}$$



- Exponential Density Function: $\{\mu\} \in \Theta$

$$f_x(x) = \frac{1}{\mu} e^{-\frac{x}{\mu}} \quad \text{for } x \geq 0$$

$$F_x(x) = 1 - e^{-\frac{x}{\mu}} \quad \text{for } x \geq 0$$

μ as avg

- ① Arithmetic avg
- ||
- ② weighted avg = expected value

- Gaussian(Normal) Density Function $\{\mu, \sigma\} \in \Theta$

$$f_x(x) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(x-\mu)^2}{2\sigma^2}}$$

Discrete Probability Functions

- Examples:

- Bernoulli Distribution:

- $$\begin{cases} 1 - p & \text{for } x = 0 \\ p & \text{for } x = 1 \end{cases}$$

In Bernoulli, just a **single** trial is conducted

- Binomial Distribution:

- $$P(X = k) = \binom{n}{k} p^k (1 - p)^{n-k}$$

k is number of successes

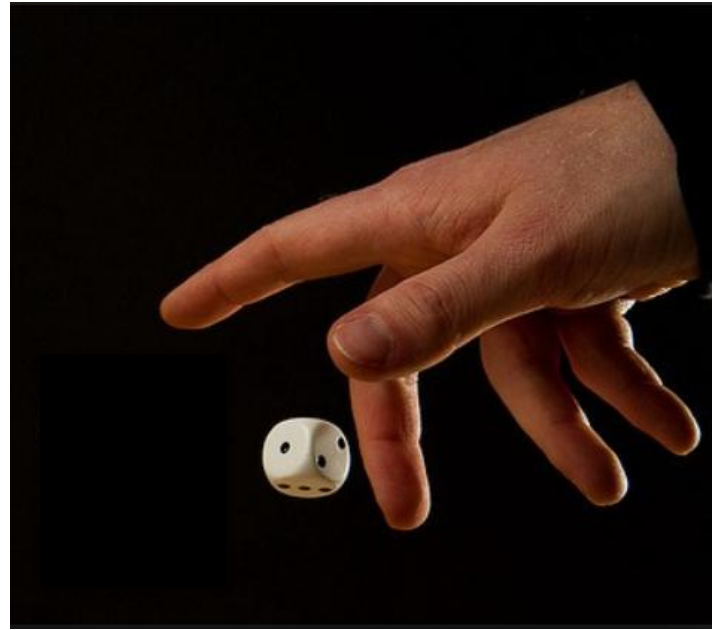
n-k is number of failures

$\binom{n}{k}$ The total number of ways of selection **k** distinct combinations of **n** trials, **irrespective of order**.

Outline

- Probability Distributions
- Joint and Conditional Probability Distributions ← 
- Bayes' Rule
- Mean and Variance
- Properties of Gaussian Distribution
- Maximum Likelihood Estimation

Example



X = Throw a
dice



Y = Flip a coin

$\mathbf{X}$ and $\mathbf{Y}$ are random variables

$\mathbf{N}$ = total number of trials

n_{ij} = Number of occurrence

		$\mathbf{X}$						
		$x_{i=1} = 1$	$x_{i=2} = 2$	$x_{i=3} = 3$	$x_{i=4} = 4$	$x_{i=5} = 5$	$x_{i=6} = 6$	C_j
$\mathbf{Y}$	$y_{j=2} = tail$	$n_{ij} = 3$	$n_{ij} = 4$	$n_{ij} = 2$	$n_{ij} = 5$	$n_{ij} = 1$	$n_{ij} = 5$	20
	$y_{j=1} = head$	$n_{ij} = 2$	$n_{ij} = 2$	$n_{ij} = 4$	$n_{ij} = 2$	$n_{ij} = 4$	$n_{ij} = 1$	15
	C_i	5	6	6	7	5	6	N=35

X**C_j**
 $x_{i=1} = 1 \quad x_{i=2} = 2 \quad x_{i=3} = 3 \quad x_{i=4} = 4 \quad x_{i=5} = 5 \quad x_{i=6} = 6$

$n_{ij} = 3$	$n_{ij} = 4$	$n_{ij} = 2$	$n_{ij} = 5$	$n_{ij} = 1$	$n_{ij} = 5$	20
$n_{ij} = 2$	$n_{ij} = 2$	$n_{ij} = 4$	$n_{ij} = 2$	$n_{ij} = 4$	$n_{ij} = 1$	15
5	6	6	7	5	6	N=35

Y $y_{j=2} = \text{tail}$ $y_{j=1} = \text{head}$ **C_i**

①

$$P(Y=t, X=5) = \frac{1}{35} = \frac{n_{ij}}{N}$$

$$\textcircled{2} \quad P(Y=t) = \frac{20}{35} = \frac{C_j}{N} \quad P(X=4) = \frac{7}{35} = \frac{C_i}{N}$$

$$\textcircled{3} \quad P(Y=t | X=4) = \frac{5}{7} = \frac{n_{ij}}{C_i} \quad P(X=5 | Y=t) = \frac{1}{20} = \frac{n_{ij}}{C_j}$$

$$P(Y=y, X=x) = \frac{n_{ij}}{N} = \left(\frac{n_{ij}}{C_j} \right) \frac{C_j}{N} = P(X=x | Y=y) P(Y=y)$$

$$= \frac{n_{ij}}{C_i} \frac{C_i}{N} = P(Y=y | X=x) P(X=x)$$

} product rule

$$P(A, B) = P(A|B) P(B) = P(B|A) P(A)$$

Probability:

$$p(X = x_i) = \frac{c_i}{N}$$

Joint probability:

$$p(X = x_i, Y = y_j) = \frac{n_{ij}}{N}$$

$p(A, B, C) = p(A|B, C) p(B, C) = p(A, B|C) p(C)$

Conditional probability:

$$p(Y = y_j | X = x_i) = \frac{n_{ij}}{c_i}$$

Sum rule

$$p(X = x_i) = \sum_{j=1}^L p(X = x_i, Y = y_j) \Rightarrow p(X) = \sum_Y P(X, Y)$$

Product rule

$$p(X = x_i, Y = y_j) = \frac{n_{ij}}{N} = \frac{n_{ij}}{c_i} \frac{c_i}{N} = p(Y = y_j | X = x_i) p(X = x_i)$$

$$p(X, Y) = p(Y|X)p(X)$$

Conditional Independence

- Examples:

$$P(H, F, V, D) = P(H|F, V, D) P(F, V, D) = P(H|F, D) P(F, V, D)$$

$$P(\text{Virus} | \text{DrinkBeer}) = P(\text{Virus})$$

iff **Virus** is independent of **Drink Beer**

$$\begin{aligned} &P(H|F, D) P(F|V, D) P(V, D) \\ &P(H|F, D) P(F|V) P(V|D) P(D) \\ &P(H|F, D) P(F|V) P(V) P(D) \end{aligned}$$

$$P(\text{Flu} | \text{Virus}, \text{DrinkBeer}) = P(\text{Flu} | \text{Virus})$$

iff **Flu** is independent of **Drink Beer**, given **Virus**

Statistical independence = independence

$$\begin{aligned} &P(\text{Headache} | \text{Flu}, \text{Virus}, \text{DrinkBeer}) \\ &= P(\text{Headache} | \text{Flu}, \text{DrinkBeer}) \end{aligned}$$

iff **Headache** is independent of **Virus**, given **Flu** and **Drink Beer**

$$\textcircled{y = x^2} \quad y = 2x$$

Assume the above independence, we obtain:

$$\begin{aligned} &P(\text{Headache}, \text{Flu}, \text{Virus}, \text{DrinkBeer}) \\ &= P(\text{Headache} | \text{Flu}, \text{Virus}, \text{DrinkBeer}) P(\text{Flu} | \text{Virus}, \text{DrinkBeer}) \\ &\quad P(\text{Virus} | \text{DrinkBeer}) P(\text{DrinkBeer}) \\ &= P(\text{Headache} | \text{Flu}, \text{DrinkBeer}) P(\text{Flu} | \text{Virus}) P(\text{Virus}) P(\text{DrinkBeer}) \end{aligned}$$

Outline

- Probability Distributions
- Joint and Conditional Probability Distributions
- Bayes' Rule 
- Mean and Variance
- Properties of Gaussian Distribution
- Maximum Likelihood Estimation

Bayes' Rule

- $P(X|Y)$ = Fraction of the worlds in which X is true given that Y is also true.

$$P(Y, X) = P(Y|X)P(X) \Rightarrow P(Y|X) = \frac{P(Y, X)}{P(X)}$$

- For example:

- H = "Having a headache"
- F = "Coming down with flu"
- $P(\text{Headache}|\text{Flu})$ = fraction of flu-inflicted worlds in which you have a headache. How to calculate?

$$P(Y|X) = \frac{P(X|Y)P(Y)}{P(X)}$$

- Definition:

$$P(X|Y) = \frac{P(X, Y)}{P(Y)} = \frac{P(Y|X)P(X)}{P(Y)}$$

Corollary:

$$P(X, Y) = P(Y|X)P(X)$$

This is called **Bayes Rule**

Bayes' Rule

$$P(a|b,c) = \frac{P(a,b,c)}{P(b,c)} = \frac{P(a,b|c) \cancel{P(c)}}{P(b|c) \cancel{P(c)}} = \frac{P(a,b|c)}{P(b|c)}$$

$$P(a|b,c)P(b|c)$$

- $$P(\text{Headache}|\text{Flu}) = \frac{P(\text{Headache}, \text{Flu})}{P(\text{Flu})}$$

$$= \frac{P(\text{Flu}|\text{Headache})P(\text{Headache})}{P(\text{Flu})}$$

Other cases:

$$P(a,b|c)$$

- $$P(Y|X) = \frac{P(X|Y)P(Y)}{P(X|Y)P(Y) + P(X|\neg Y)P(\neg Y)}$$

$$P(a,b) = P(a|b,c)P(b|c)$$

- $$P(Y = y_i|X) = \frac{P(X|Y)P(Y)}{\sum_{i \in S} P(X|Y = y_i)P(Y = y_i)}$$

- $$P(Y|X, Z) = \frac{P(X|Y, Z)P(Y, Z)}{P(X, Z)}$$

$$= \frac{P(X|Y, Z)P(Y, Z)}{P(X|Y, Z)P(Y, Z) + P(X|\neg Y, Z)P(\neg Y, Z)}$$

Outline

- Probability Distributions
- Joint and Conditional Probability Distributions
- Bayes' Rule
- Mean and Variance 
- Properties of Gaussian Distribution
- Maximum Likelihood Estimation

$E[\cdot]$ $\leadsto$ expected value

Mean and Variance

$E(\cdot)$ $\leadsto$ function

- Expectation: The mean value, center of mass, first moment:

$$E_X[g(X)] = \int_{-\infty}^{\infty} g(x)p_X(x)dx = \mu$$

$$E[g(x)] = \sum_{i=1}^N p(x^{(i)}) g(x^{(i)})$$

$g(x) = x$

- N-th moment: $g(x) = x^n$

- N-th central moment: $g(x) = (x - \mu)^n$

- Mean: $E_X[X] = \int_{-\infty}^{\infty} xp_X(x)dx$

- $E[\alpha X] = \alpha E[X]$

- $E[\alpha + X] = \alpha + E[X]$

$$E[\alpha] = \alpha$$

$$E[(1,2,3)] = 2$$

$$E[E[(1,2,3)]] = E[\quad]$$

$$Var(x) = \frac{1}{N} \sum_{i=1}^N (x^{(i)} - \mu)^2$$

- Variance(Second central moment): $Var(x) =$

$$E_X[(X - E_X[X])^2] = E_X[X^2] - E_X[X]^2$$

$$E[(x - E[x])^2] \checkmark$$

- $Var(\alpha X) = \alpha^2 Var(X)$

- $Var(\alpha + X) = Var(X)$

Mean and average

$$X = [1, 2, 3]$$

$$g(x) = x$$

$$p(x) = \left[\frac{1}{6}, \frac{3}{6}, \frac{2}{6}\right] \leadsto \text{pmf}$$

$$E[g(x)] = \sum_{i=1}^N p(x) g(x)$$

$$E[g(x)] = \frac{1}{6} * 1 + \frac{3}{6} * 2 + \frac{2}{6} * 3 = \frac{13}{6}$$

$$\text{Avg}(x) = \frac{1}{N} \sum_{i=1}^N g(x) = \frac{1+2+3}{3} = 2$$

$$\text{Avg}(x) \neq \frac{13}{6}$$

$$X = [1, 2, 2, 2, 3, 3] = \frac{1+2+2+2+3+3}{6} = \frac{13}{6}$$

Variance and average:

$h = \text{height} = x_1$

$$X = \begin{bmatrix} 1 \\ 2 \\ 3 \end{bmatrix}_{n \times d=1}$$

$x_1^{\{23\}}$

$$\mu_h = \frac{1+2+3}{3} = 2$$

$$\sigma_h^2 = \frac{1}{N} \sum_{i=1}^N (x_1^{(i)} - \mu_h)^2$$

$$\frac{1}{N} \sum (.) = \text{avg} (.) = E[.]$$

$$E[(x - \mu_h)^2]$$

$$E[(x_i - E[x_i])^2]$$

$X \rightarrow \bar{X}$ → centering

$$\bar{X} = \begin{bmatrix} \bar{h} = h - \mu_h \\ 1 - \mu_h \\ 2 - \mu_h \\ 3 - \mu_h \end{bmatrix}_{n \times d} = \begin{bmatrix} \bar{h} \\ -1 \\ 0 \\ 1 \end{bmatrix}$$

$\mu_{\bar{h}} = 0$

$$\text{Var}(x) = \frac{1}{N} \bar{X}_{d \times n}^T \bar{X}_{n \times d} = \frac{1}{N} \begin{bmatrix} 1 - \mu_h & 2 - \mu_h & 3 - \mu_h \end{bmatrix} \begin{bmatrix} 1 - \mu_h \\ 2 - \mu_h \\ 3 - \mu_h \end{bmatrix} = \frac{1}{N} \left((1 - \mu_h)^2 + (2 - \mu_h)^2 + (3 - \mu_h)^2 \right)$$

Covariance:

$$X = \begin{bmatrix} h & w \\ 1 & 4 \\ 2 & 5 \\ 3 & 6 \end{bmatrix}$$

$$\mu_h = \frac{1+2+3}{3} = 2 \quad \mu_w = \frac{4+5+6}{3} = 5$$

$$\sigma_h^2 = \frac{1}{N} \sum_{i=1}^N (x_1^{(i)} - \mu_h)^2$$

$$\sigma_w^2 = \frac{1}{N} \sum_{i=1}^N (x_2^{(i)} - \mu_w)^2$$

$$X \rightsquigarrow \bar{X} = \begin{bmatrix} \bar{h} & \bar{w} \\ 1-\mu_h & 4-\mu_w \\ 2-\mu_h & 5-\mu_w \\ 3-\mu_h & 6-\mu_w \end{bmatrix}$$

$$\mu_{\bar{h}} = 0 \quad \mu_{\bar{w}} = 0$$

$$\Rightarrow \text{COV} = \left(\frac{1}{N} \bar{X}^T \bar{X} \right)_{d \times d}$$

Symmetric

$$\text{COV} = \frac{1}{N}$$

$$\begin{bmatrix} 1-\mu_h & 2-\mu_h & 3-\mu_h \\ 4-\mu_w & 5-\mu_w & 6-\mu_w \end{bmatrix}$$

$$\begin{bmatrix} 1-\mu_h & 4-\mu_w \\ 2-\mu_h & 5-\mu_w \\ 3-\mu_h & 6-\mu_w \end{bmatrix}$$

=

$$\begin{bmatrix} h & w \\ \sigma_h^2 & \sigma_{hw} \\ \sigma_{wh} & \sigma_w^2 \end{bmatrix}$$

$$\sigma_{hw} = \sigma_{wh}$$

Correlation:

Standardization $\rightarrow$ Standard deviation $\rightarrow \sigma$

$$X \rightarrow \bar{X} \rightarrow \bar{X}^* = \begin{bmatrix} \frac{h - \mu_h}{\sigma_h} & \frac{w - \mu_w}{\sigma_w} \\ \frac{1 - \mu_h}{\sigma_h} & \frac{2 - \mu_w}{\sigma_w} \\ \vdots & \vdots \end{bmatrix}$$

$$\text{COR} = \frac{1}{N} \bar{X}^{*T} \bar{X}^* = h \begin{bmatrix} \frac{\sigma_h^2}{\sigma_h^2} = 1 & \sigma_{hw} \\ -\sigma_{hw} & \frac{\sigma_w^2}{\sigma_w^2} = 1 \end{bmatrix} \quad \text{EDA}$$

$$y = x^2$$

$$\sigma_{hw} = 1 \quad h = 2w \quad \sigma_{hw} = 0 \quad h \perp w$$

$$\sigma_{hw} = 0.99 \quad \left(h = 2\underline{w} + \alpha \underline{E}^{\text{noise}} \right)$$

For Joint Distributions

$$\text{Cov}(X, Y) = \cancel{E[Z^3]} - \cancel{E[Z]} E[Z^2] = 0$$

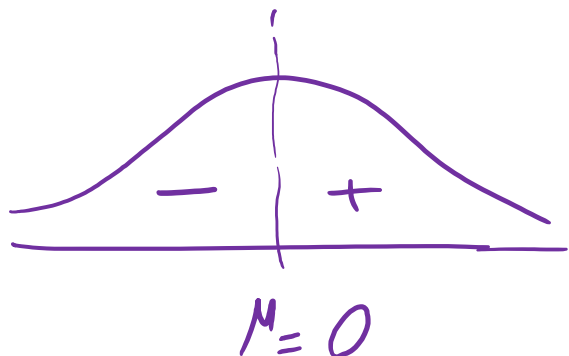
$\downarrow$ 0 $\downarrow$ 0

• Expectation and Covariance:

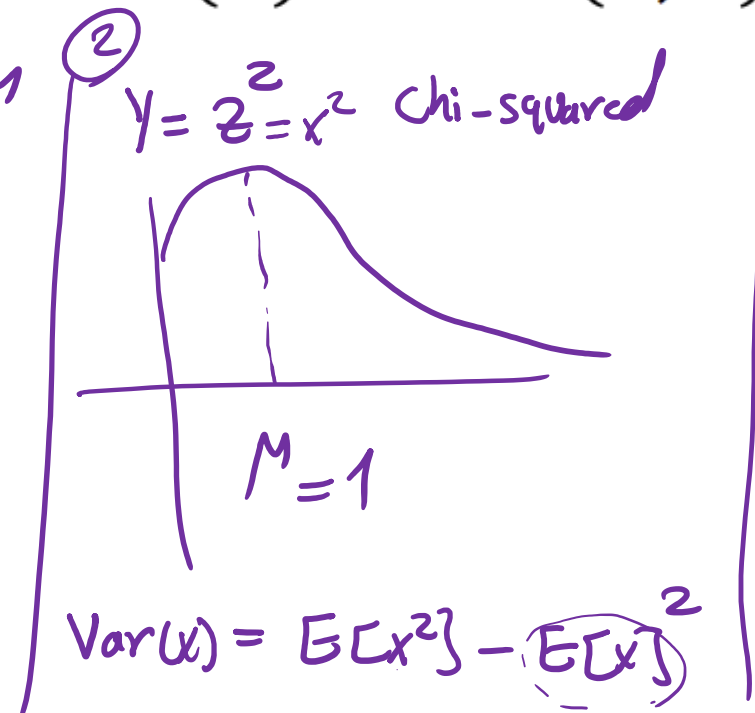
- $E[X + Y] = E[X] + E[Y]$
- $\text{cov}(X, Y) = E[(X - E_X[X])(Y - E_Y[Y])] = \underline{E[XY] - E[X]E[Y]}$
- $\text{Var}(X + Y) = \text{Var}(X) + 2\text{cov}(X, Y) + \text{Var}(Y)$

① Z standard gaussian distribution

$\downarrow$
 $\mu = 0$ $\sigma = 1$
 $X = Z$



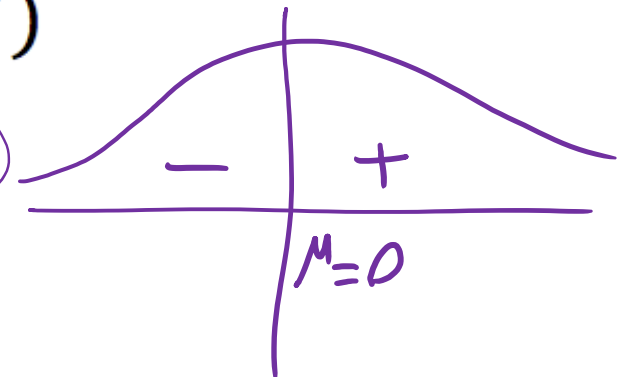
② $Y = Z^2 = X^2$ Chi-squared



$$\text{Var}(X) = E[X^2] - \cancel{(E[X])^2}$$

$$1 = E[Y] - 0 \Rightarrow E[Y] = 1$$

$g = Z^3$



Outline

- Probability Distributions
- Joint and Conditional Probability Distributions
- Bayes' Rule
- Mean and Variance
- Properties of Gaussian Distribution ← 
- Maximum Likelihood Estimation

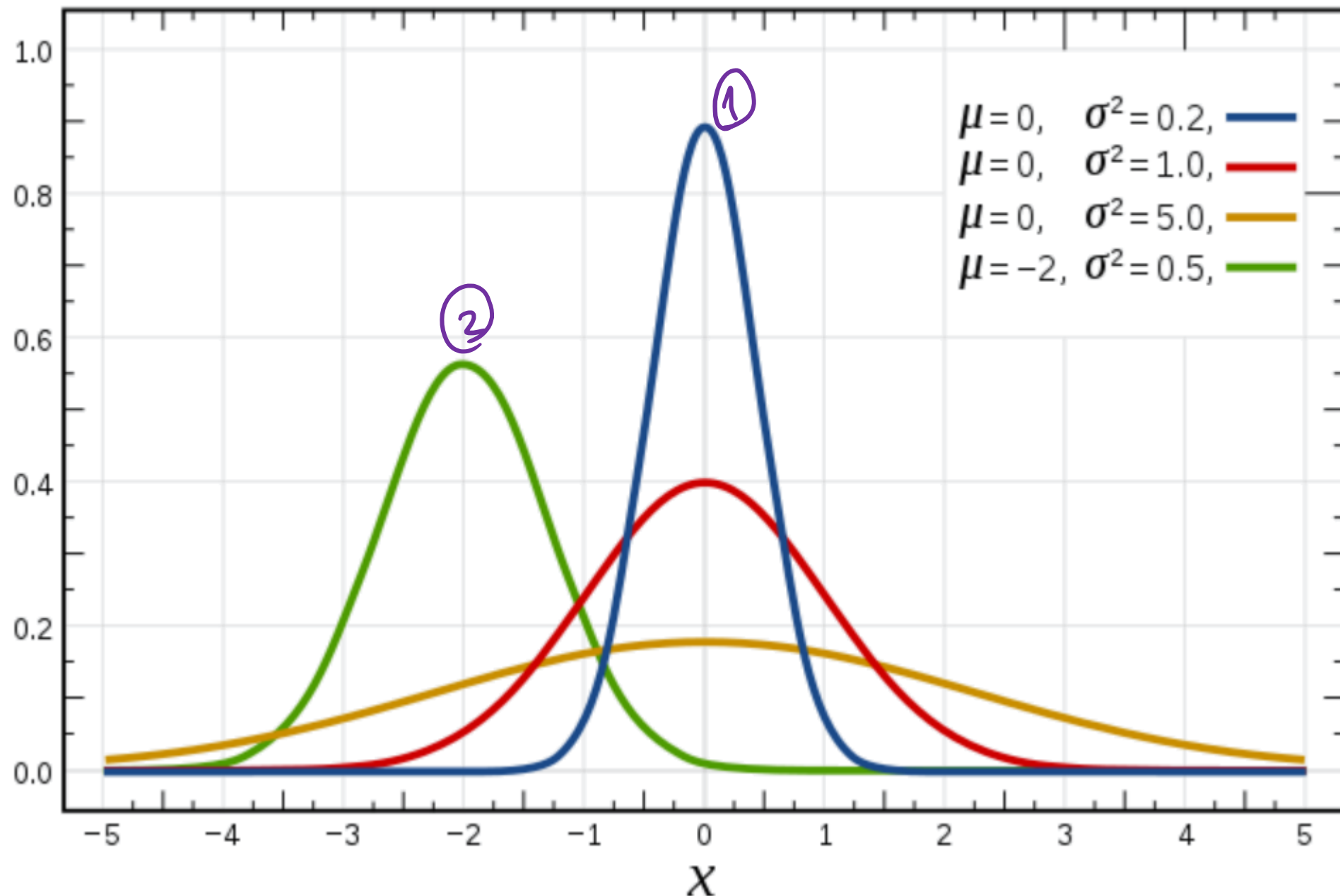
Gaussian Distribution

Univariate

- Gaussian Distribution:

$$f(x|\mu, \sigma^2) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(x-\mu)^2}{2\sigma^2}}$$

Probability density function



Probability versus likelihood

Prob vs Likelihood

$$P(X=h) = \frac{1}{2} \Rightarrow H, H, H, T \Rightarrow P(X=H) = \frac{3}{4}$$

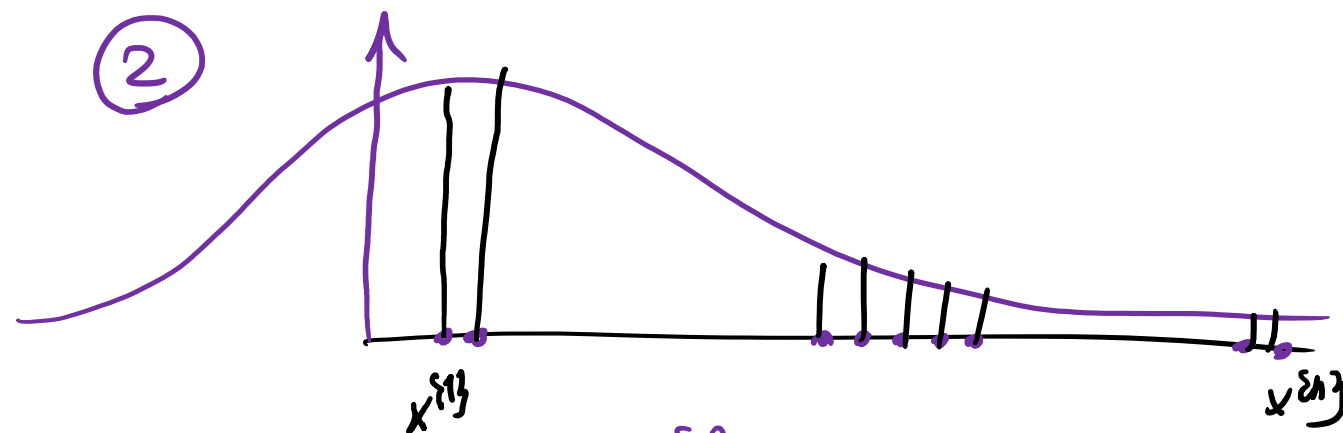
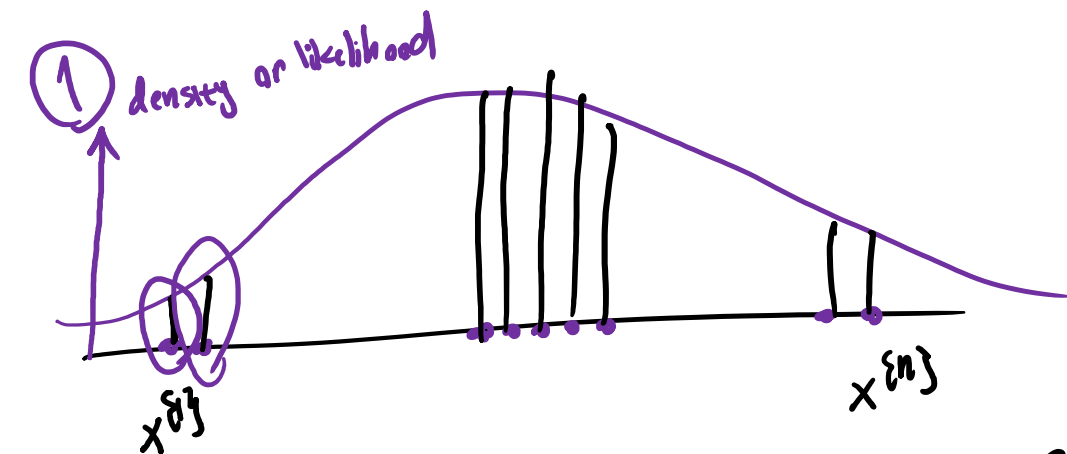
probability

Statistic world

most Lazy form Lazy form

$$\underline{f(x)} = f(x|\Theta) = f(x|a,b) = \frac{1}{\sqrt{2\pi a^2}} \exp\left(-\frac{(x-b)^2}{2a^2}\right)$$

function



Objective function $\overset{\text{Max}}{L(a,b|x)}$

$\rightarrow a=1, b=1 \Rightarrow f(x^{(1)}, x^{(2)}, \dots, x^{(n)} | a=1, b=1) =$

$\rightarrow a=1, b=2 \Rightarrow f(x^{(1)}, \dots, x^{(n)} | a=1, b=2) = f(x^{(1)} | a,b) f(x^{(2)} | a,b) \dots f(x^{(n)} | a,b)$

Prob vs Likelihood

→ Likelihood function

$$\log(a \cdot b) = \log a + \log b$$

$$L(a, b | x) = \prod_{i=1}^N f(x^{(i)} | a, b)$$

Objective function

max

$$l(a, b | x) = \log L(a, b | x) = \sum_{i=1}^N \log f(x^{(i)} | a, b)$$

$$\frac{\partial l(a, b | x)}{\partial a} = 0$$

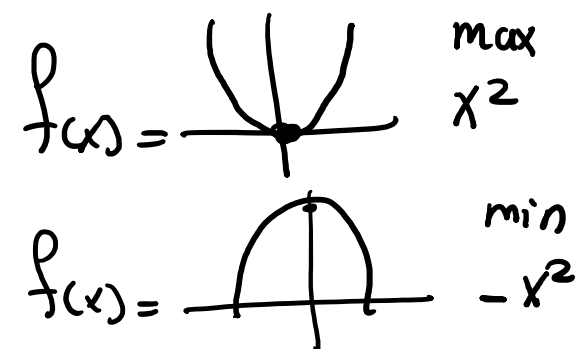
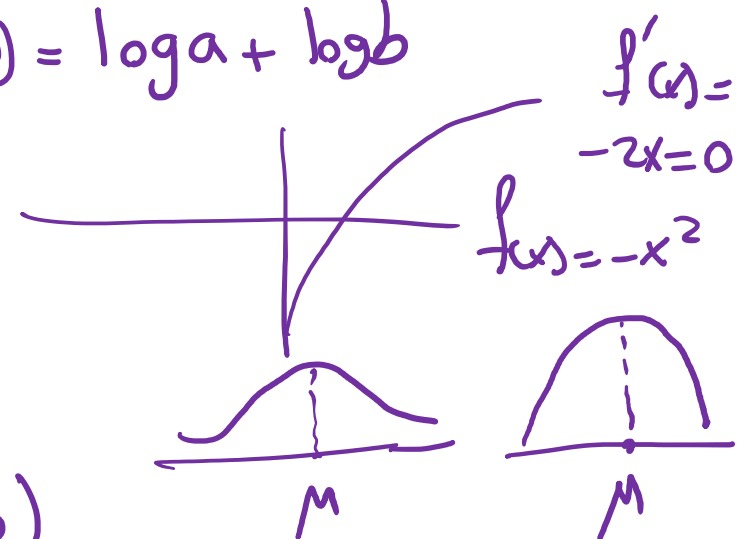
↓

$$a = \sigma^2$$

$$\frac{\partial l(a, b | x)}{\partial b} = 0$$

$$b = \mu = \frac{1}{N} \sum_{i=1}^N x^{(i)}$$

MLE



Multivariate Gaussian Distribution

$X = \begin{bmatrix} \text{h} & \text{w} \end{bmatrix}$

$$\Sigma_{d \times d} = \frac{1}{N} \bar{X}^T \bar{X} = \frac{1}{N} \begin{bmatrix} \text{h} & \text{w} \\ \text{w} & \text{h} \end{bmatrix} \begin{bmatrix} \sigma_h^2 & \sigma_{hw} \\ \sigma_{wh} & \sigma_w^2 \end{bmatrix}$$

$$p(x|\mu, \Sigma) = \frac{1}{(2\pi)^{n/2} |\Sigma|^{1/2}} \exp\left\{-\frac{1}{2} (x - \mu)^T \Sigma^{-1} (x - \mu)\right\}$$

$\mu = [\mu_h \ \mu_w]_{1 \times d}$
 behaves like a datapoint
 $|\Sigma|$ covariance
 μ 1xd

- Moment Parameterization $\mu = E(X)$

$$\Sigma = \text{Cov}(X) = E[(X - \mu)(X - \mu)^T]$$

- Mahalanobis Distance $\Delta^2 = (x - \mu)^T \Sigma^{-1} (x - \mu)$
- Tons of applications (MoG, FA, PPCA, Kalman filter,...)

Properties of Gaussian Distribution

- The **linear transform** of a Gaussian r.v. is a Gaussian. Remember that no matter how x is distributed

$$E(AX + b) = AE(X) + b$$

$$\text{Cov}(AX + b) = A\text{Cov}(X)A^T$$

this means that for Gaussian distributed quantities:

$z \propto x$

$$X \sim N(\mu, \Sigma) \rightarrow Z = AX + b \sim N(A\mu + b, A\Sigma A^T)$$

μ_{new} Σ_{new}

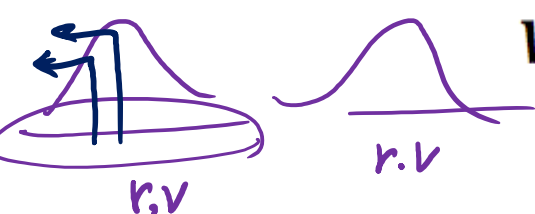
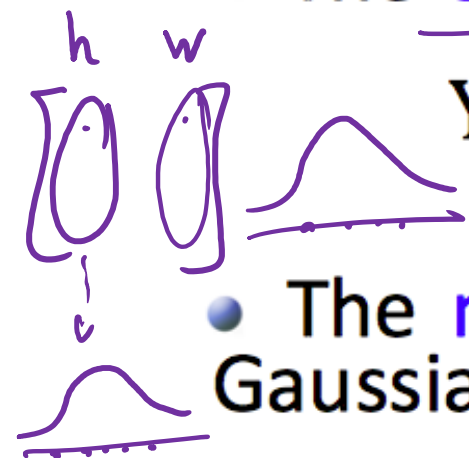
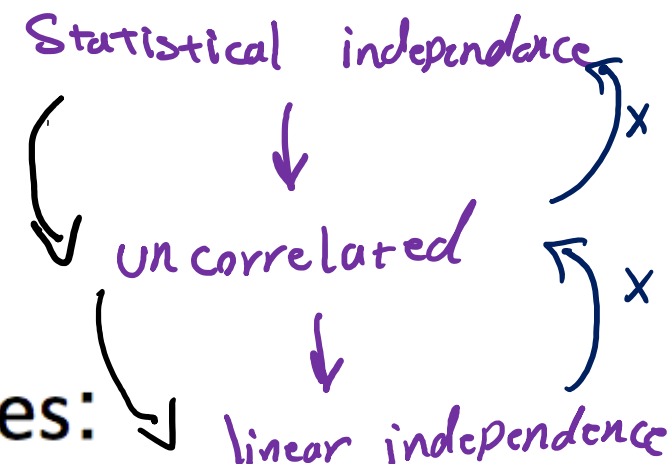
- The **sum** of two independent Gaussian r.v. is a Gaussian

$$Y = X_1 + X_2, X_1 \perp X_2 \rightarrow \mu_y = \mu_1 + \mu_2, \Sigma_y = \Sigma_1 + \Sigma_2$$

- The **multiplication** of two Gaussian functions is another Gaussian function (although no longer normalized)

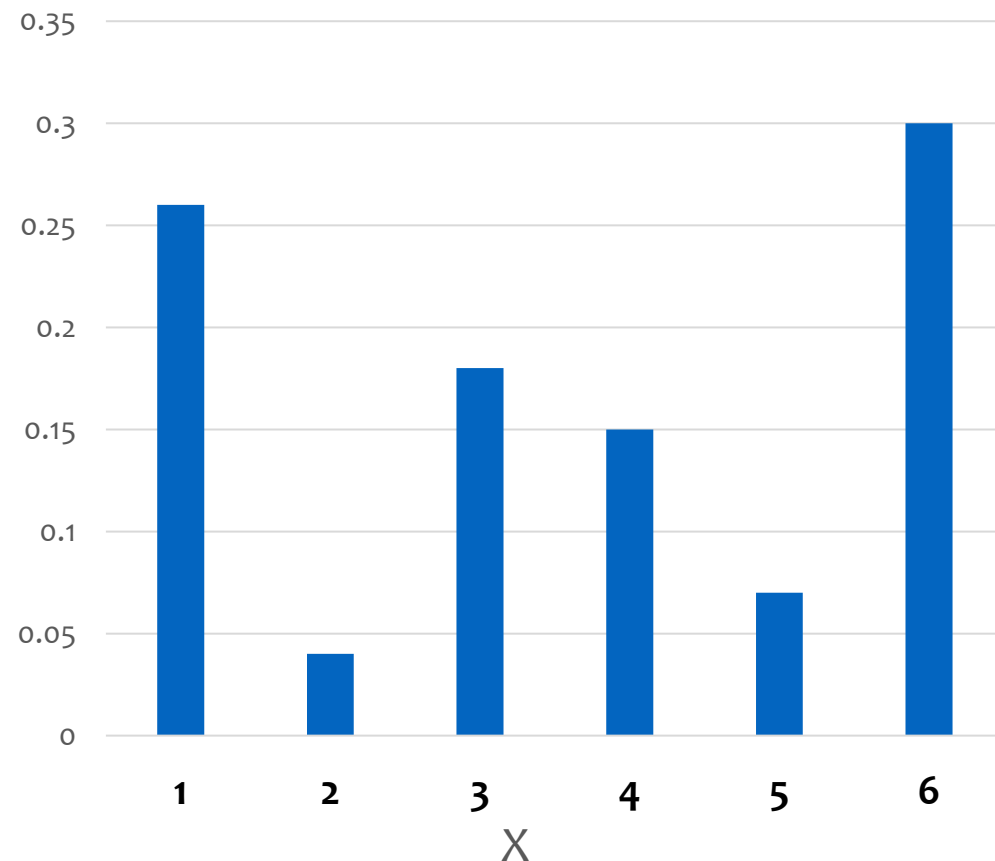
$$N(a, A)N(b, B) \propto N(c, C),$$

$$\text{where } C = (A^{-1} + B^{-1})^{-1}, c = CA^{-1}a + CB^{-1}b$$



Central Limit Theorem CLT

Probability mass function of a **biased** dice



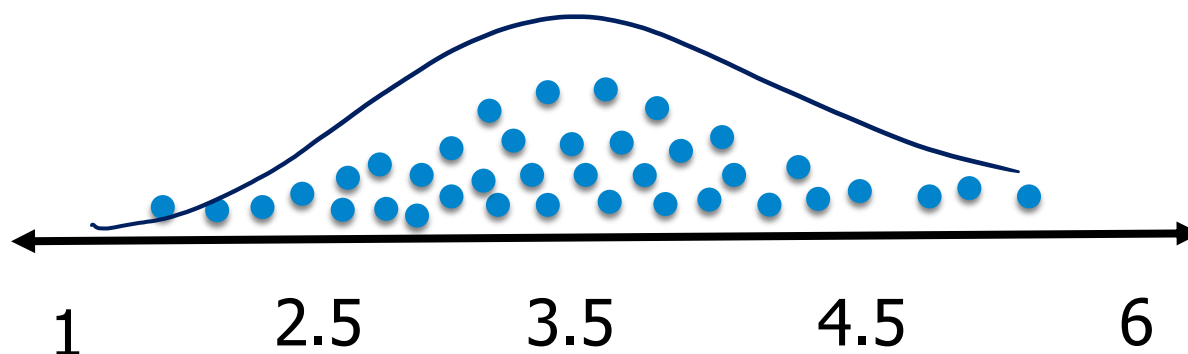
Let's say, I am going to get a sample from this pmf having a size of **$n = 4$**

$$S_1 = \{1, 1, 1, 6\} \Rightarrow E(S_1) = 2.25$$

$$S_2 = \{1, 1, 3, 6\} \Rightarrow E(S_2) = 2.75$$

$\vdots$

$$S_m = \{1, 4, 6, 6\} \Rightarrow E(S_m) = 4.25$$



According to CLT, it will follow a bell curve distribution (normal distribution)

Outline

- Probability Distributions
- Joint and Conditional Probability Distributions
- Bayes' Rule
- Mean and Variance
- Properties of Gaussian Distribution
- Maximum Likelihood Estimation 

Maximum Likelihood Estimation

- Probability: inferring probabilistic quantities for data given fixed models (e.g. prob. of events, marginals, conditionals, etc).
- Statistics: inferring a model given fixed data observations (e.g. clustering, classification, regression).

Main assumption:

Independent and identically distributed random variables
i.i.d

$$P(x^{i3}) = \underline{\underline{\theta}}^{x^{i3}} (1-\theta)^{(1-x^{i3})} \leadsto \text{pmf} = P(x^{i3} | \theta)$$

$$\Rightarrow \text{Objective function } P(x^{\{1\}}, x^{\{2\}}, \dots, x^{\{N\}} | \theta) = P(x^{\{1\}} | \theta) P(x^{\{2\}} | \theta) \dots P(x^{\{N\}} | \theta)$$

$$\text{Max} \quad L(\theta | x) = \prod_{i=1}^N P(x^{i3} | \theta)$$

$$l(\theta | x) = \log L(\theta | x) = \sum_{i=1}^N \log P(x^{i3} | \theta) = \sum_{i=1}^N \log (\theta^{x^{i3}} (1-\theta)^{(1-x^{i3})})$$

Maximum Likelihood Estimation

For Bernoulli (i.e. flip a coin):

Objective function: $P(x_i|\theta) = \theta^{x^{\{i\}}} (1 - \theta)^{1-x^{\{i\}}}$ $x^{\{i\}} \in \{0,1\}$ or *head, tail*

$$L(\theta|X) = L(\theta|X = x^{\{1\}}, X = x^{\{2\}}, X = x^{\{3\}}, \dots, X = x^{\{n\}})$$

i.i.d assumption

$$L(\theta|X) = \prod_{i=1}^n P(x^{\{i\}}|\theta)$$

$$L(\theta|X) = \prod_{i=1}^n P(x^{\{i\}}|\theta) = \prod_{i=1}^n \theta^{x^{\{i\}}} (1 - \theta)^{1-x^{\{i\}}}$$

$$\begin{aligned} L(\theta|X) &= \theta^{x^{\{1\}}} (1 - \theta)^{1-x^{\{1\}}} \times \theta^{x^{\{2\}}} (1 - \theta)^{1-x^{\{2\}}} \dots \times \theta^{x^{\{n\}}} (1 - \theta)^{1-x^{\{n\}}} = \\ &= \theta^{\sum x^{\{i\}}} (1 - \theta)^{\sum (1-x^{\{i\}})} \end{aligned}$$

We don't like multiplication, let's convert it into summation

What's the trick?

Take the log

$$L(\theta|X) = \overbrace{\theta^{\sum x^{(i)}}}^a \overbrace{(1-\theta)^{\sum (1-x^{(i)})}}^b$$

$$\log \theta^{\sum x^{(i)}} = \sum x^{(i)} \log \theta$$

$\log a + \log b$

$$\log L(\theta|X) = l(\theta|X) = \log(\theta) \sum_{i=1}^n x^{(i)} + \log(1-\theta) \sum_{i=1}^n (1-x^{(i)})$$

How to optimize θ ?

$$\frac{\partial l(\theta|X)}{\partial \theta} = 0 \quad \frac{\sum_{i=1}^n x^{(i)}}{\theta} - \frac{\sum_{i=1}^n (1-x^{(i)})}{1-\theta} = 0$$

$$\theta = \frac{1}{n} \sum_{i=1}^n x^{(i)}$$

$x^{\text{head}} = 1$
 $x^{\text{tail}} = 0$

H, H, H, T, T

$$\theta = \frac{1+1+1+0+0}{5} = \frac{3}{5}$$