



10-315 Introduction to ML

MLE and Probabilistic Formulation of Machine Learning

Instructor: Pat Virtue

Logistics

Exam

- Wed in class, see Piazza for details

Coming this week

- Probability Primer
- MAP Pre-reading
- Mini project info

Today

- Wrap-up regularization (for now)
- MLE
 - Maximum likelihood estimation
 - Probabilistic formulation of linear and logistic regression

$$f_X(x) = \text{Gaussian}$$

$$\begin{aligned} &\rightarrow P(x) \\ &\rightarrow P(X=x) \end{aligned}$$

$$X \text{ event} \rightarrow \text{num} \\ P(Y=\text{dog})$$

Poll 1

Course feedback link on Piazza

Are you finished with the course feedback form?

- A. Yes
- B. No
- C. Still working on it

Exercises

Calculate the probability of these event sequences happening

1. Coin

a) Fair:

$\{H, H, T, H\}$

b) Biased, $\phi = 3/4$ heads

$\{H, H, T, H\}$

2. 4-sided die with sides: A, B, C, D

a) Fair:

$\{A, B, D, D, A\}$

b) Weighted, $[\phi_A, \phi_B, \phi_C, \phi_D] = [1/10, 2/10, 3/10, 4/10]$

$\{A, B, D, D, A\}$

Exercises

Calculate the probability of these event sequences happening

1. Coin

a) Fair:

{H, H, T, H}

$$\frac{1}{2} \cdot \frac{1}{2} \cdot \frac{1}{2} \cdot \frac{1}{2} = \frac{1}{16}$$

b) Biased, $\phi = 3/4$ heads

{H, H, T, H}

$$\phi \cdot \phi \cdot (1-\phi) \cdot \phi = \frac{3}{4} \cdot \frac{3}{4} \cdot \frac{1}{4} \cdot \frac{3}{4} = \frac{27}{256}$$

2. 4-sided die with sides: A, B, C, D

a) Fair:

{A, B, D, D, A}

$$\frac{1}{6} \cdot \frac{1}{6} \cdot \frac{1}{6} \cdot \frac{1}{6} \cdot \frac{1}{6}$$

b) Weighted, $[\phi_A, \phi_B, \phi_C, \phi_D] = [1/10, 2/10, 3/10, 4/10]$

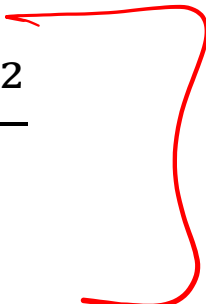
{A, B, D, D, A}

$$\phi_A \cdot \phi_B \cdot \phi_D \cdot \phi_D \cdot \phi_A = \frac{1}{10} \cdot \frac{2}{10} \cdot \frac{4}{10} \cdot \frac{4}{10} \cdot \frac{1}{10}$$

Poll 2

Implement a function in Python for the pdf of a Gaussian distribution.

Python numpy or math packages are fine, no scipy, etc.

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


```
def gaussian(x, mu, sigmaSq):
```

What is `gaussian(3.3, 2.2, 1.1)`?

Poll 3

0.0000023

→ 0.0000401

Assume that exam scores are drawn independently from the same Gaussian (Normal) distribution.

Given three exam scores {75, 80, 90}, which pair of parameters is a better fit (a higher likelihood)?

A) Mean 80, standard deviation 3

☒ B) Mean 85, standard deviation 7

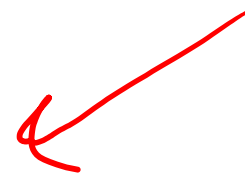
C) I don't know

$p(75|80, 3) p(80|80, 3) p(90|80, 3)$

$p(75|85, 7) p(80|85, 7) p(90|85, 7)$

Use a calculator/computer.

Gaussian PDF: $p(y | \mu, \sigma^2) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(y-\mu)^2}{2\sigma^2}}$



Poll 4

Trick coin: comes up heads only $\frac{1}{3}$ of the time

1 flip: H probability: $\frac{1}{3}$

2 flips: H,H probability: $\frac{1}{3} \cdot \frac{1}{3}$

3 flips: H,H,T probability: $\frac{1}{3} \cdot \frac{1}{3} \cdot \left(1 - \frac{1}{3}\right)$

But what property allows us to just multiply these?

indep.

Poll 5

$$\mathcal{D} = \{y^{(1)}, y^{(2)}, y^{(3)}\}$$

Which of these is the correct labeling of these properties?

$$Y \sim \text{Bern}(\phi)$$

ident.
distr.

product
rule

indep.

$$p(\underline{\mathcal{D}} | \theta) = p(Y_1 = y^{(1)}, Y_2 = y^{(2)}, Y_3 = y^{(3)} | \phi_1, \phi_2, \phi_3)$$

$$= p(Y_1 = y^{(1)}, Y_2 = y^{(2)}, Y_3 = y^{(3)} | \phi)$$

$$= p(Y_1 | \phi) p(Y_2 | \cancel{Y_1}, \phi) p(Y_3 | \cancel{Y_1}, \cancel{Y_2}, \phi)$$

$$= p(Y_1 = y^{(1)} | \phi) p(Y_2 = y^{(2)} | \phi) p(Y_3 = y^{(3)} | \phi)$$

$$= \prod_{i=1}^N p(Y = y^{(i)} | \phi)$$

Likelihood

Pre-reading

Likelihood

Likelihood: The probability (or density) of random variable Y taking on value y given the distribution parameters, θ .

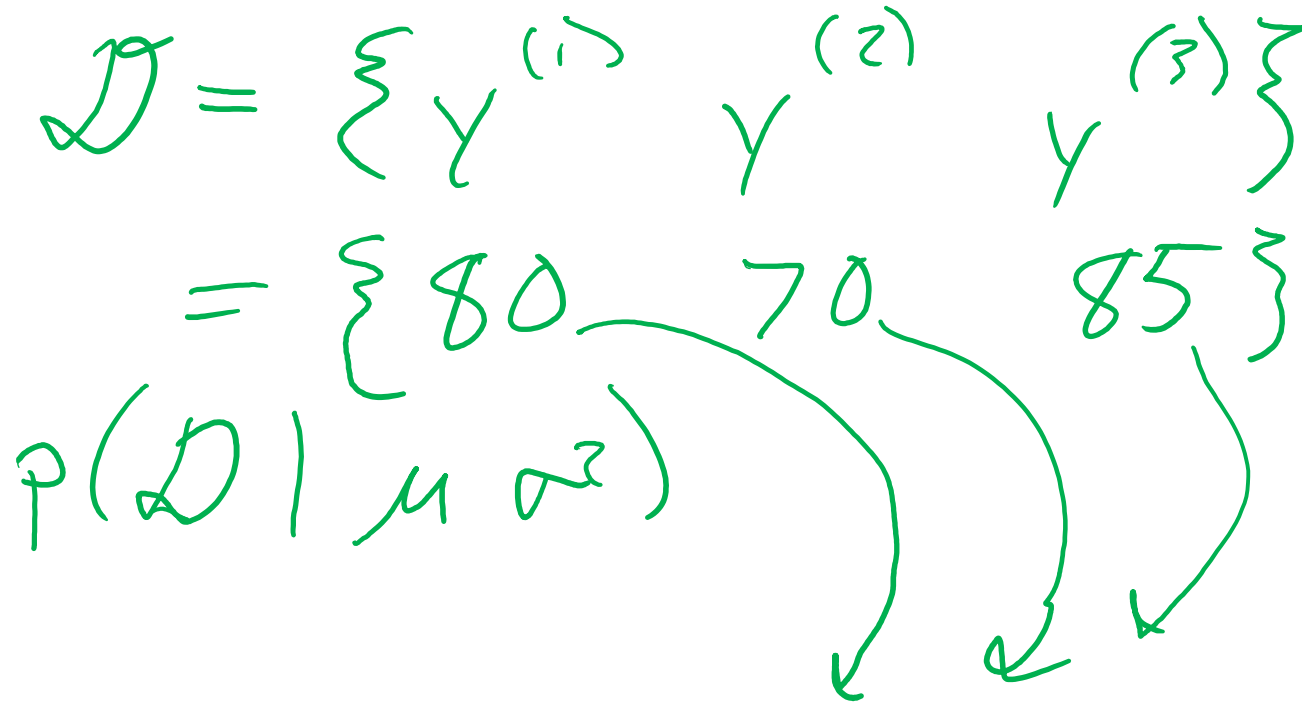
Likelihood

Likelihood: The probability (or density) of random variable Y taking on value y given the distribution parameters, θ .

Grades

$$\begin{aligned} \mathcal{D} &= \{y^{(1)} \quad y^{(2)} \quad y^{(3)}\} \\ &= \{80 \quad 70 \quad 85\} \end{aligned}$$

$p(\mathcal{D} \mid \mu, \sigma^2)$



 Gaussian PDF: $p(y \mid \mu, \sigma^2) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(y-\mu)^2}{2\sigma^2}}$

Likelihood

$$\phi = \frac{1}{3}$$

Trick coin: comes up heads only $\frac{1}{3}$ of the time

1 flip: H	probability: $\frac{1}{3}$
2 flips: H,H	probability: $\frac{1}{3} \cdot \frac{1}{3}$
3 flips: H,H,T	probability: $\frac{1}{3} \cdot \frac{1}{3} \cdot \left(1 - \frac{1}{3}\right)$

But why can we just multiply these?

Handwritten green ink showing the likelihood function $p(D | \theta)$ and a specific instance $p(HHT | \phi)$ circled with an arrow pointing to the general form.

Likelihood and i.i.d

Likelihood: The probability (or density) of random variable Y taking on value y given the distribution parameters, θ .

i.i.d.: Independent and identically distributed

$$p(D | \theta)$$

ident $\rightarrow p(Y_1 = y^{(1)}, Y_2 = y^{(2)}, Y_3 = y^{(3)} | \phi_1, \phi_2, \phi_3)$

$$p(Y = y^{(1)}, Y = y^{(2)}, Y = y^{(3)} | \phi)$$

indep. $\rightarrow p(Y = y^{(1)} | \phi) p(Y = y^{(2)} | \phi) p(Y = y^{(3)} | \phi)$

$$\prod_{i=1}^n p(Y = y^{(i)} | \phi)$$

Bernoulli Likelihood

Bernoulli distribution:

$$Y \sim \text{Bern}(\phi) \quad p(y \mid \phi) = \begin{cases} \phi, & y = 1 \\ 1 - \phi, & y = 0 \end{cases}$$

What is the likelihood for three i.i.d. samples, given parameter ϕ :

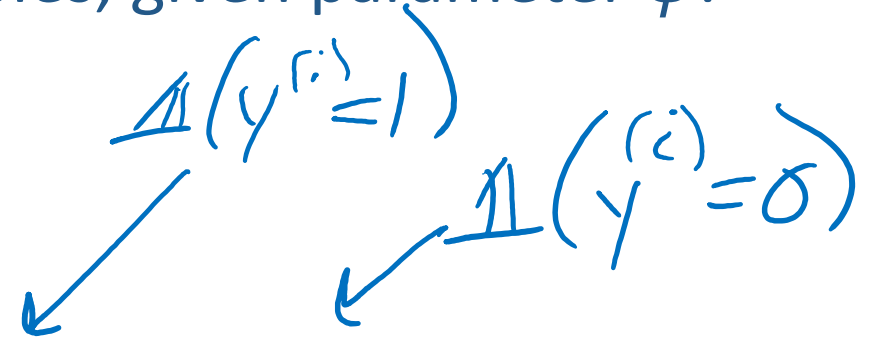
$$\mathcal{D} = \{y^{(1)} = 1, y^{(2)} = 1, y^{(3)} = 0\}$$

$$\prod_{i=1}^N p(Y = y^{(i)} \mid \phi)$$

$$= \underline{\phi} \cdot \underline{\phi} \cdot (1 - \phi)$$

$$\underline{\phi^1 (1 - \phi)^0} \cdot \underline{\phi^1 (1 - \phi)^0} \cdot \underline{\phi^0 (1 - \phi)^1}$$

$\mathbb{1}(y^{(i)} = 1)$ $\mathbb{1}(y^{(i)} = 0)$

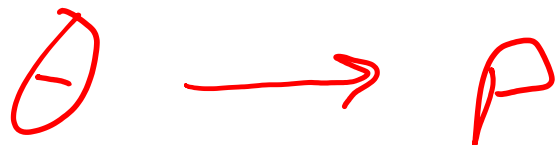


MLE

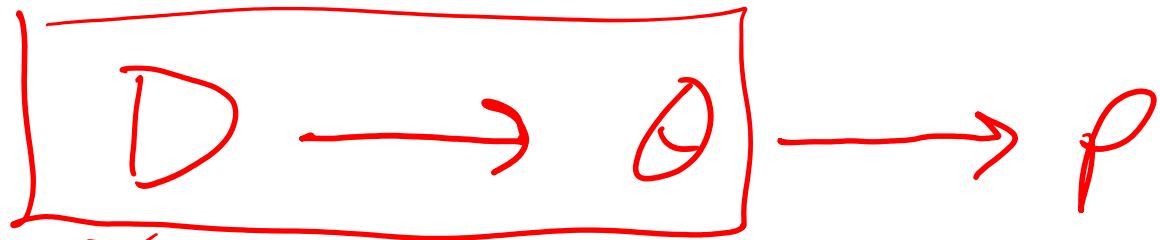
Maximum likelihood estimation

From Probability to Statistics

Prob



Stat



x
 x
 x

predict

generate
 x_{new}

Estimating Parameters with Likelihood

We model the outcome of a single mysterious weighted-coin flip as a Bernoulli random variable:

$$Y \sim \text{Bern}(\phi)$$
$$p(y \mid \phi) = \begin{cases} \phi, & y = 1 \text{ (heads)} \\ 1 - \phi, & y = 0 \text{ (tails)} \end{cases}$$

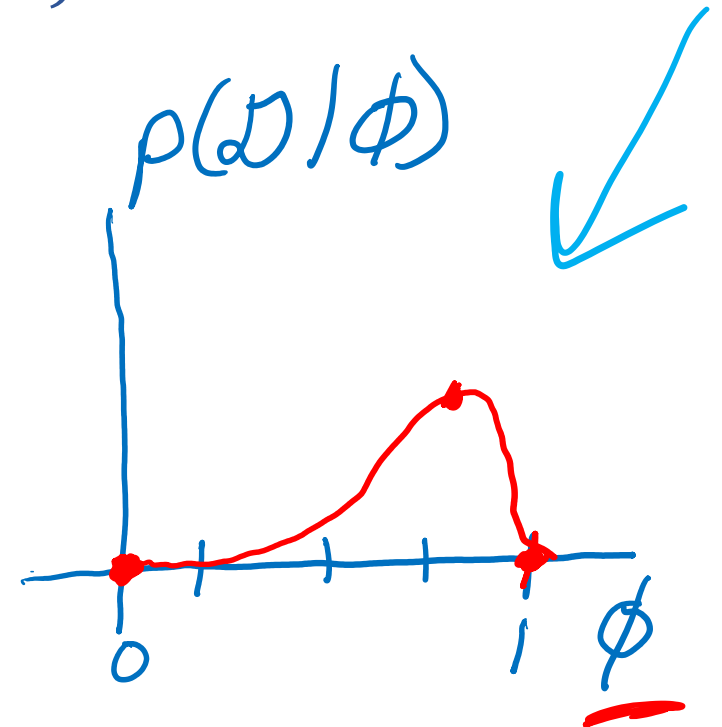
$\uparrow$

Given the ordered sequence of coin flip outcomes:
[1, 0, 1, 1]

What is the estimate of parameter $\hat{\phi}$?

$$p(D \mid \phi) = \phi \cdot \phi \cdot (1 - \phi) \cdot \phi$$
$$= \phi^3 (1 - \phi)^1$$

$\leftarrow 3/4$



Likelihood and Maximum Likelihood Estimation

Likelihood: The probability (or density) of random variable Y taking on value y given the distribution parameters, θ .

$$p(D|\theta) = \prod_{i=1}^n p(y^{(i)}|\theta)$$

Likelihood function: The value of likelihood as we change theta (same as likelihood, but conceptually we are considering many different values of the parameters)

$$L(\theta; D) = p(D|\theta)$$

Maximum Likelihood Estimation (MLE): Find the parameter value that maximizes the likelihood.

$$\hat{\theta}_{MLE} = \underset{\theta}{\operatorname{argmax}} L(\theta)$$

MLE as Data Increases

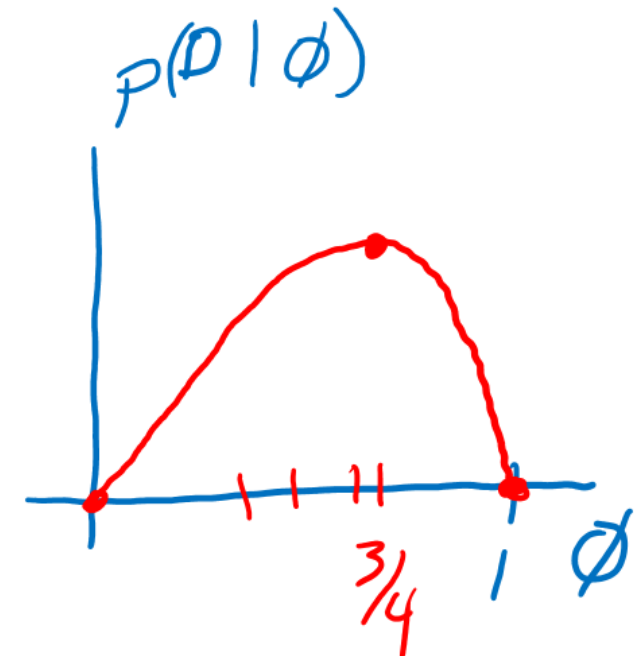
Given the ordered sequence of coin flip outcomes:
[1, 0, 1, 1]

$$p(\mathcal{D} \mid \phi) = \prod_i^N p(y^{(i)} \mid \phi) = \phi^{N_{y=1}} (1 - \phi)^{N_{y=0}}$$

$\phi (1 - \phi) \phi \phi$

What happens as we flip more coins?

better MLE with more data



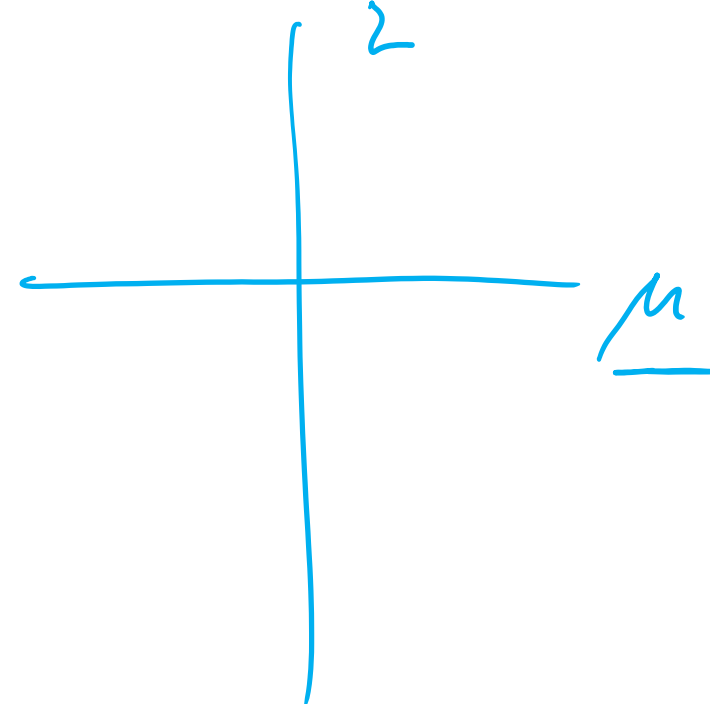
MLE for Gaussian

Gaussian distribution:

$$Y \sim \mathcal{N}(\mu, \sigma^2)$$

$$p(y \mid \mu, \sigma^2) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(y-\mu)^2}{2\sigma^2}}$$

$$\mathcal{D} = \{y^{(1)} = 65, y^{(2)} = 95, y^{(3)} = 85\}$$



Formulate the likelihood for three i.i.d. samples, given parameters μ, σ^2 ?

$$L(\mu, \sigma^2) = \prod_{i=1}^N \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(y^{(i)}-\mu)^2}{2\sigma^2}}$$

$$= \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(65-\mu)^2}{2\sigma^2}} \cdot \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(95-\mu)^2}{2\sigma^2}} \cdot \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(85-\mu)^2}{2\sigma^2}}$$

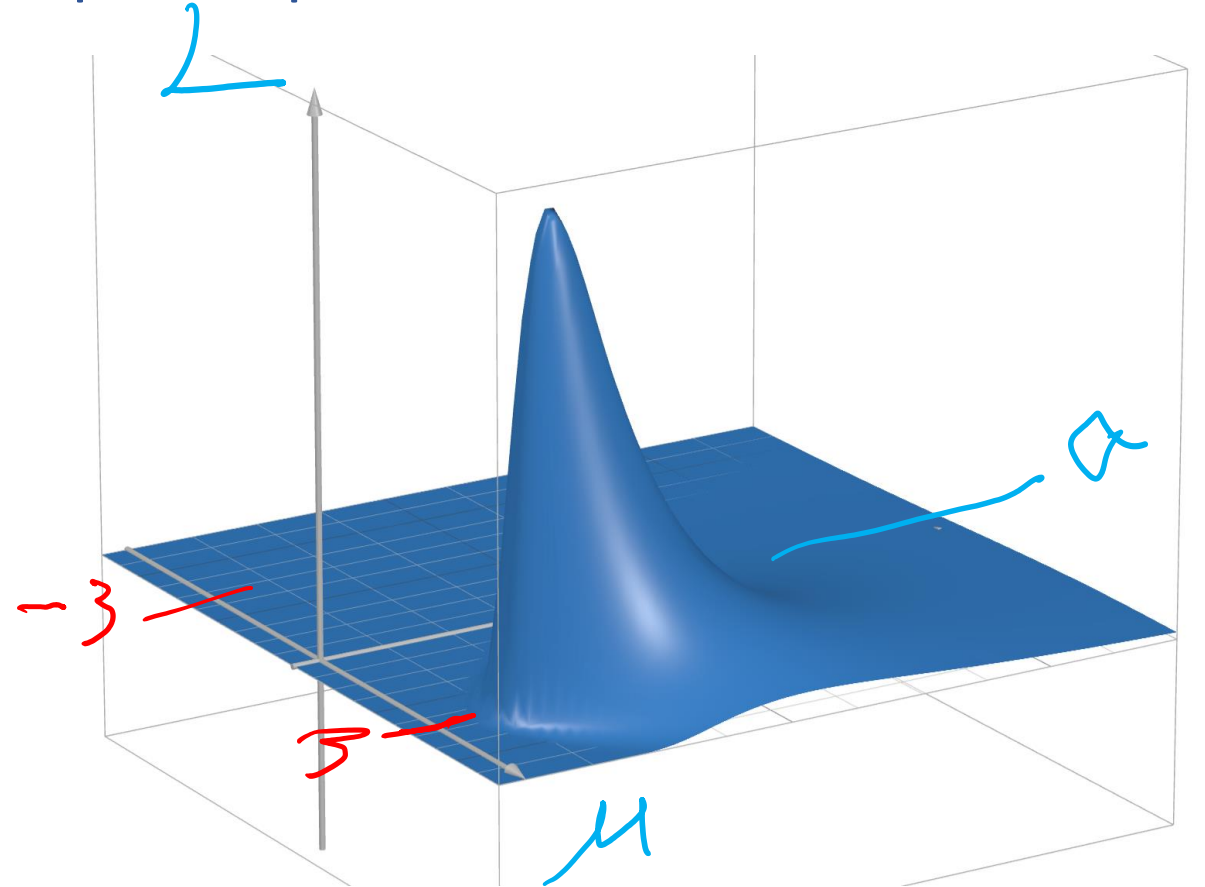
$$\hat{\theta}_{MLE} = \operatorname{argmax}_{\theta} \prod_i p(y^{(i)} \mid \theta)$$

MLE for Gaussian

Assume that exam scores are drawn independently from the same Gaussian (Normal) distribution.

Given three exam scores 2, 3, 4, which pair of parameters is the best fit (the highest likelihood)?

$$p(\mathcal{D}|\mu, \sigma^2) = \prod_{i=1}^N \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(y^{(i)} - \mu)^2}{2\sigma^2}}$$



MLE

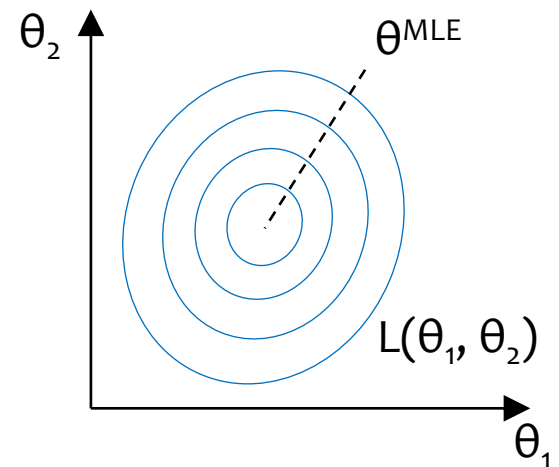
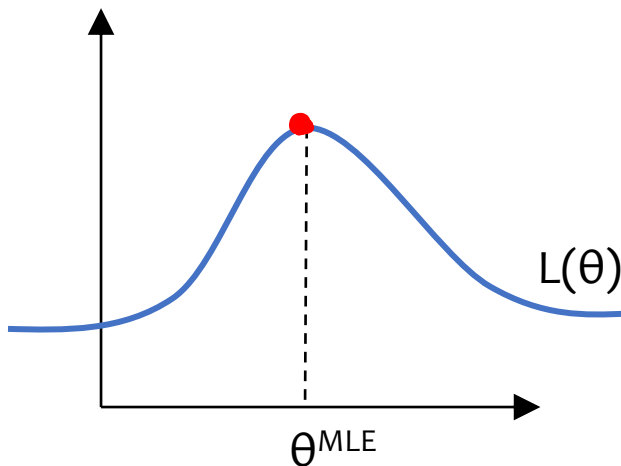
Suppose we have data $\mathcal{D} = \{x^{(i)}\}_{i=1}^N$

Principle of Maximum Likelihood Estimation:

Choose the parameters that maximize the likelihood of the data.

$$\theta^{\text{MLE}} = \underset{\theta}{\operatorname{argmax}} \prod_{i=1}^N p(\mathbf{x}^{(i)} | \theta)$$

Maximum Likelihood Estimate (MLE)



$$\log mn = \log m + \log n$$

Likelihood and Log Likelihood

Likelihood: The probability (or density) of random variable Y taking on value y given the distribution parameters, θ .

$$p(\mathcal{D} \mid \theta) = \prod_{i=1}^N p(y^{(i)} \mid \theta)$$

Likelihood function: The value of likelihood as we change theta (same as likelihood, but conceptually we are considering many different values of the parameters)

$$\mathcal{L}(\theta; \mathcal{D}) = \prod p(y^i \mid \theta)$$

$$l(\theta; \mathcal{D}) = \log \mathcal{L}(\theta; \mathcal{D}) = \sum_{i=1}^N \log p(y^i \mid \theta)$$

MLE Objective

Updating our objective: Minimize neg. log likelihood

$$\begin{aligned}\hat{\theta}_{MLE} &= \operatorname{argmax}_{\theta} \prod_{i=1}^N p(y^{(i)} | \theta) \\ &= \operatorname{argmax}_{\theta} \sum_{i=1}^N \log p(y^{(i)} | \theta) \\ &= \operatorname{argmin}_{\theta} - \sum_{i=1}^N \log p(y^{(i)} | \theta)\end{aligned}$$

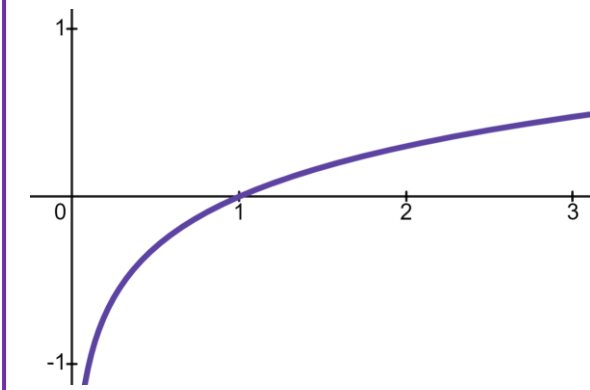
Minimize $J(\theta) = -\log \mathcal{L}(\theta; \mathcal{D}) = - \sum_{i=1}^N \log p(y^{(i)} | \theta)$

Benefits
of log

Log is monotonic:

If $a < b$

then $\log a < \log b$



Easier math
 $\pi \rightarrow \Sigma$

$e^z \rightarrow z$
calculus

Numerical stability

MLE for Gaussian

Gaussian distribution:

$$Y \sim \mathcal{N}(\mu, \sigma^2)$$

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

$$\begin{aligned} & p(75|80,3) \quad p(80|80,3) \quad p(90|80,3) \\ & 0.033 \times 0.133 \times 0.0005 = 0.0000023 \\ & \log \hookrightarrow -3.4 - 2.0 - 7.6 = -13.0 \end{aligned}$$

What is the log likelihood for three i.i.d. samples, given parameters μ, σ^2 ?

$$\mathcal{D} = \{y^{(1)} = 75, y^{(2)} = 80, y^{(3)} = 90\}$$

$$L(\mu, \sigma^2) = \prod_{i=1}^N \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(y^{(i)}-\mu)^2}{2\sigma^2}}$$

$$\ell(\mu, \sigma^2) = \sum_{i=1}^N -\log \sqrt{2\pi\sigma^2} - \frac{(y^{(i)} - \mu)^2}{2\sigma^2}$$

$$\hat{\theta}_{MLE} = \operatorname{argmax}_{\theta} \prod_i p(y^{(i)} | \theta)$$

$$\hat{\theta}_{MLE} = \operatorname{argmax}_{\theta} \sum_i \log p(y^{(i)} | \theta)$$

Recipe for Estimation

MLE

1. Formulate the likelihood, $p(\mathcal{D} \mid \theta)$
 2. Set objective $J(\theta)$ equal to negative log of likelihood
$$J(\theta) = -\log p(\mathcal{D} \mid \theta)$$
 3. Compute derivative of objective, $\partial J / \partial \theta$
 4. Find $\hat{\theta}$, either
 - a. Set derivate equal to zero and solve for θ
 - b. Use (stochastic) gradient descent to step towards better θ
- 

Probabilistic Formulation for ML

MLE for Linear and Logistic Regression

Using Statistics for Machine Learning

Likelihood vs conditional likelihood

(conditional) likelihood

$$\prod p(y^i | x^i, \theta)$$

likelihood

$$p(D | \theta)$$
$$= \prod_{i=1}^N p(x^i, y^i | \theta)$$

$$D = \{x^i, y^i\}$$

Recipe for Estimation

$$\pi_p(y|x, \theta)$$


MLE

1. Formulate the likelihood, $p(\mathcal{D} \mid \theta)$
2. Set objective $J(\theta)$ equal to negative log of likelihood
$$J(\theta) = -\log p(\mathcal{D} \mid \theta)$$
3. Compute derivative of objective, $\partial J / \partial \theta$
4. Find $\hat{\theta}$, either
 - a. Set derivate equal to zero and solve for θ
 - b. Use (stochastic) gradient descent to step towards better θ

$$J(\theta) = \sum (-y^i (\log \hat{y}^i) - (1 - \hat{y}^i) \log(1 - \hat{y}^i))$$

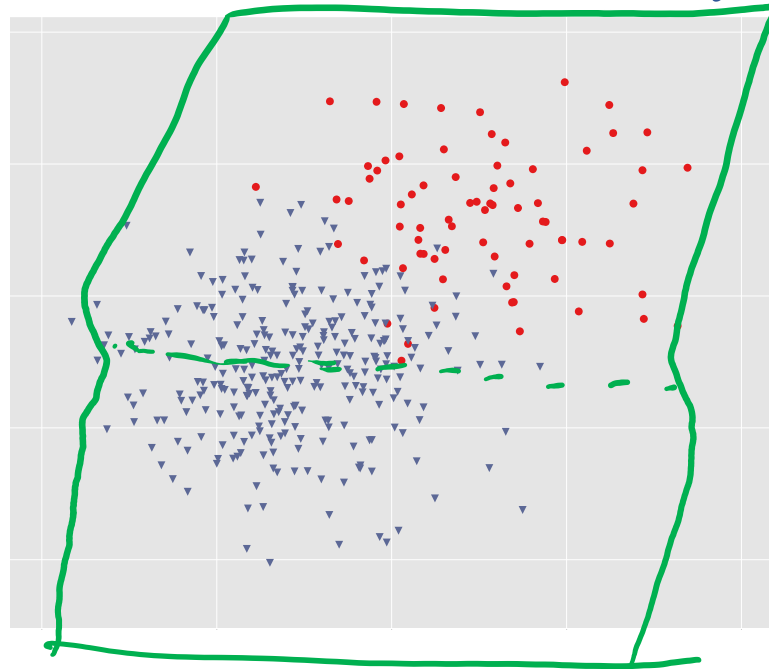
M(C)LE for Logistic Regression

$$p(y^{(i)} = 1 | x^{(i)}; w, b)$$

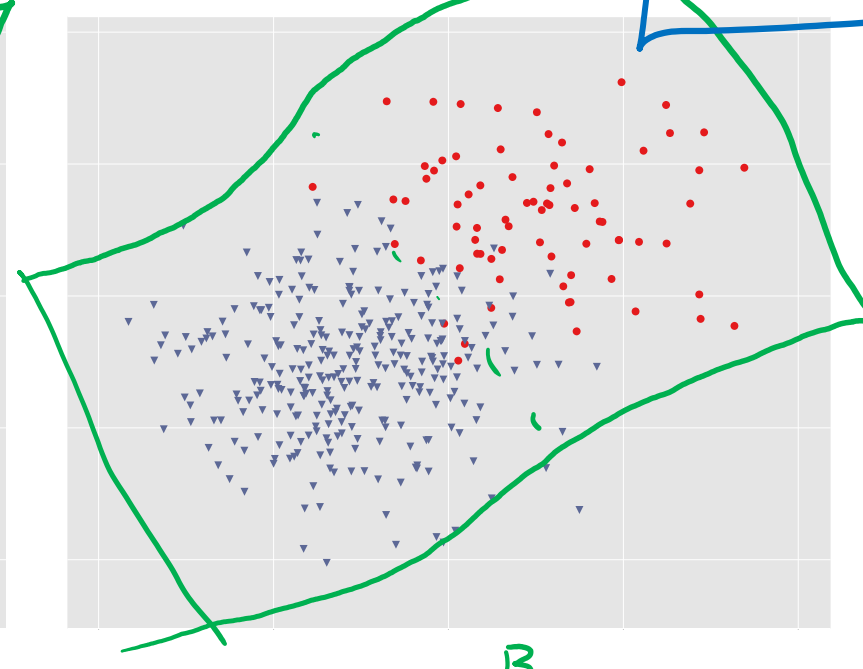
$\uparrow$ $\hat{y}$

Learn to predict if a patient has cancer ($Y = 1$) or not ($Y = 0$) given the input of just one test results, X_A and X_B

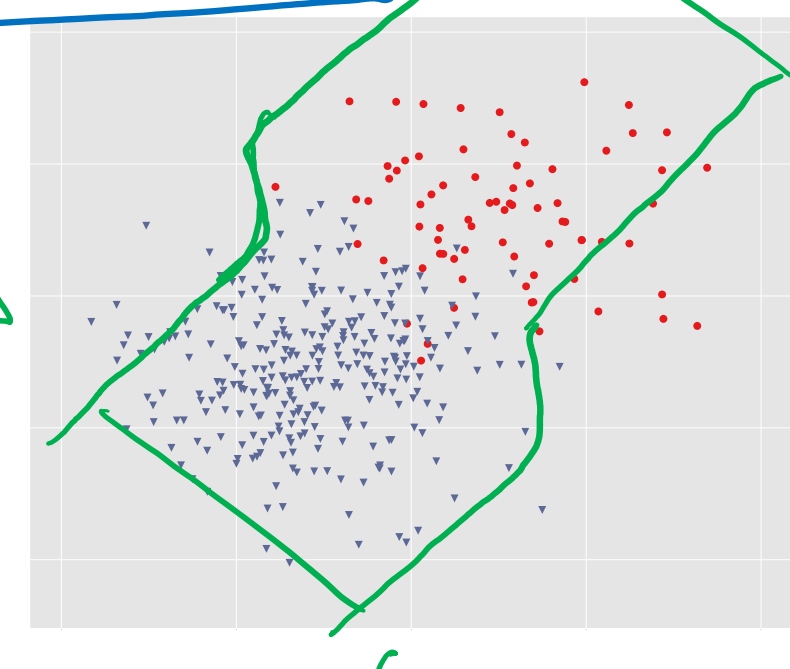
$$\hat{\theta}_{MLE} = \operatorname{argmax}_{\theta} \prod_i^N \frac{1}{1 + e^{-\theta^T x^{(i)}}}^{\mathbb{I}(y^{(i)}=1)} \left(1 - \frac{1}{1 + e^{-\theta^T x^{(i)}}} \right)^{\mathbb{I}(y^{(i)}=0)}$$



θ^A



θ^B

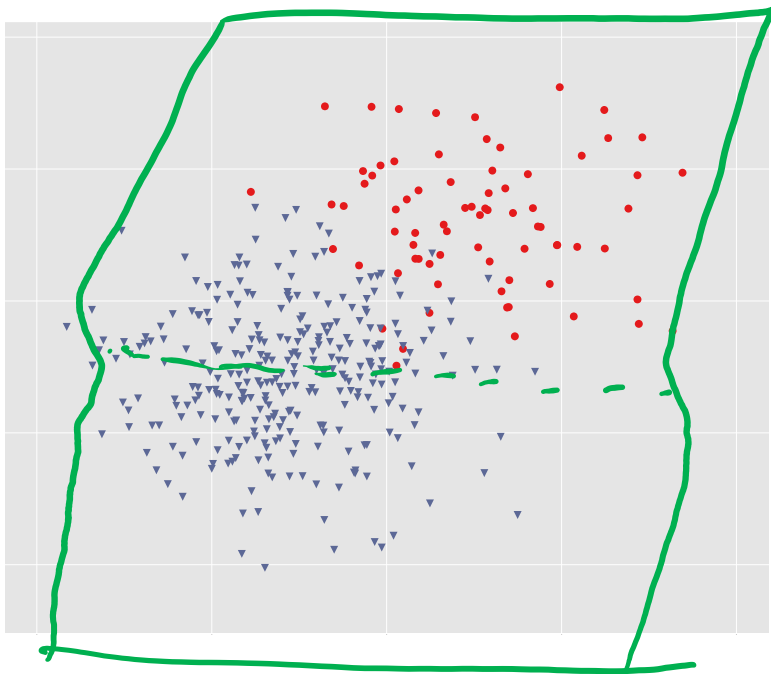


θ^C

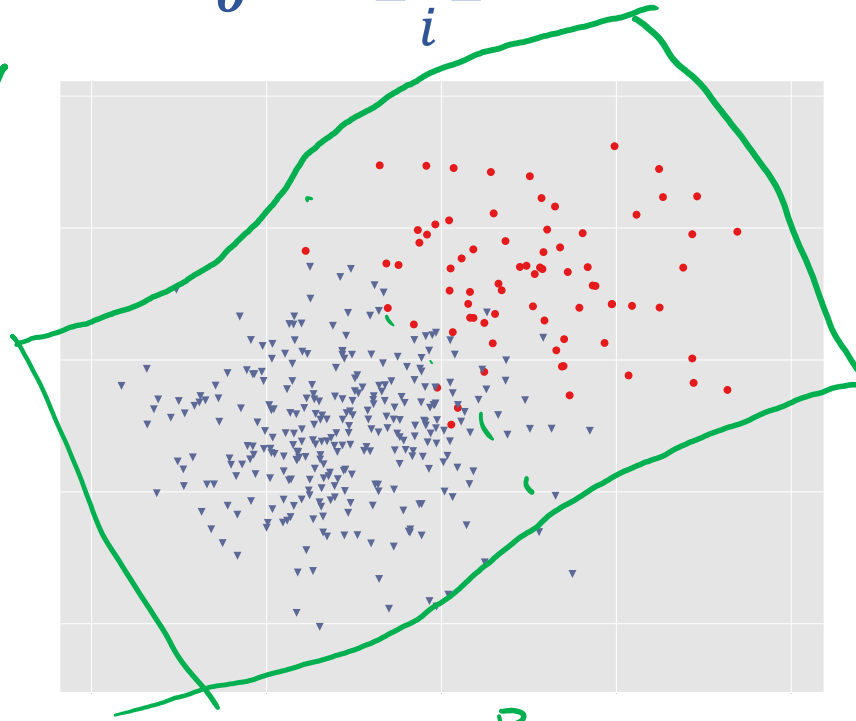
M(C)LE for Logistic Regression

Learn to predict if a patient has cancer ($Y = 1$) or not ($Y = 0$) given the input of just one test results, X_A and X_B

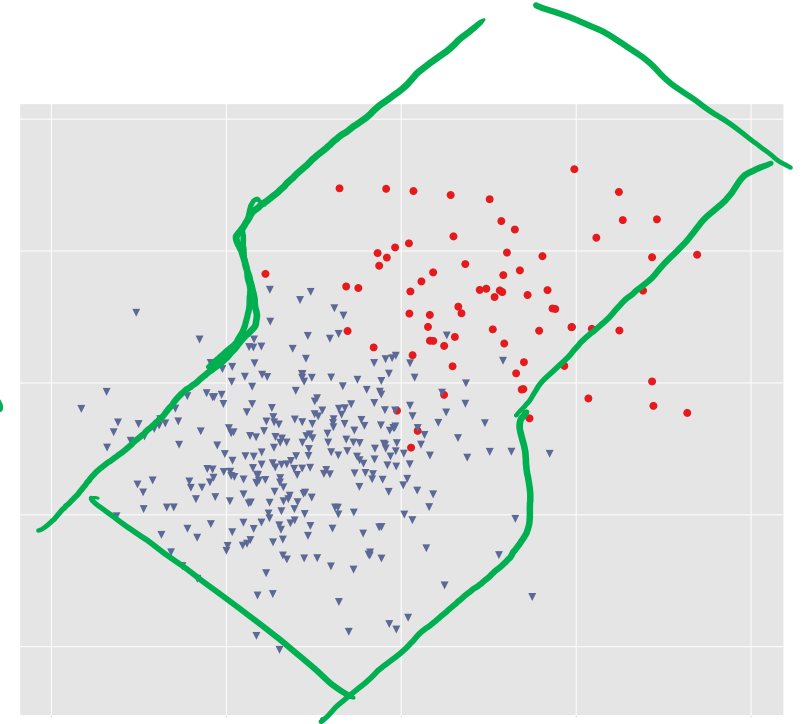
$$\hat{\theta}_{MLE} = \operatorname{argmax}_{\theta} \prod_i p(y^{(i)} \mid \mathbf{x}^{(i)}, \theta)$$



θ^A



θ^B



θ^C

M(C)LE for Multi-class Logistic Regression

Learn to predict if probability of output belonging to class k , Y_k , given input X , $P(Y_k = 1 \mid X, \boldsymbol{\theta}_1, \dots, \boldsymbol{\theta}_K)$

$$\hat{\Theta}_{MLE} = \underset{\Theta}{\operatorname{argmax}} \prod_i^N \prod_k^K \frac{e^{\boldsymbol{\theta}_k^T \mathbf{x}^{(i)}} \mathbb{I}(y_k^{(i)} = 1)}{\sum_{l=1}^K e^{\boldsymbol{\theta}_l^T \mathbf{x}^{(i)}}}$$



M(C)LE for Multi-class Logistic Regression

Learn to predict if probability of output belonging to class k , Y_k , given input X , $P(Y_k = 1 \mid X, \boldsymbol{\theta}_1, \dots, \boldsymbol{\theta}_K)$

$$\mathcal{L}(\Theta; \mathcal{D}) = \prod_i^N \prod_k^K \frac{e^{\boldsymbol{\theta}_k^T \mathbf{x}^{(i)}} \mathbb{I}(y_k^{(i)} = 1)}{\sum_{l=1}^K e^{\boldsymbol{\theta}_l^T \mathbf{x}^{(i)}}}$$

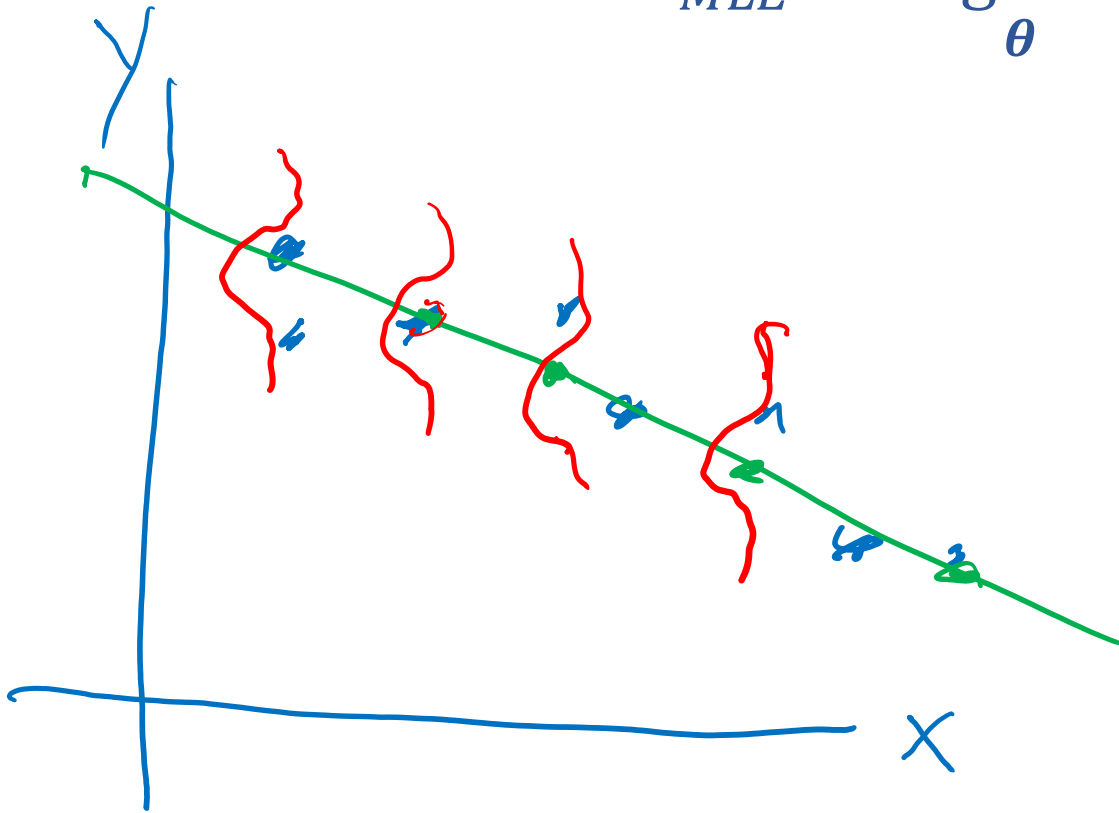
$$J(\theta) = -\log \mathcal{L}(\theta; \mathcal{D})$$

M(C)LE for Linear Regression

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

Probabilistic interpretation of linear regression

$$\hat{\theta}_{MLE} = \operatorname{argmax}_{\theta} \prod_i^N p(y^{(i)} | \mathbf{x}^{(i)}, \theta)$$



$$\begin{aligned} y &= \underline{w^T x + b} + \epsilon \\ \epsilon &\sim \mathcal{N}(0, \sigma^2) \\ \rightarrow y &\sim \mathcal{N}(\underline{w^T x + b}, \sigma^2) \end{aligned}$$

$\theta^T \tilde{x}$

M(C)LE for Linear Regression

Probabilistic interpretation of linear regression

$$y \sim \mathcal{N}(\theta^T x, \sigma^2)$$

$y^{(i)}$

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

$\theta^T x$

$$\mathcal{L}(\theta; \mathcal{D}) = \prod_i p(y^{(i)} | \mathbf{x}^{(i)}, \theta)$$

$\mu = \theta^T x$

$$= \prod_i \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(y^{(i)} - \theta^T x^{(i)})^2}{2\sigma^2}}$$

$$\prod_{i=1}^N \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(z^{(i)} - \mu)^2}{2\sigma^2}}$$

$$\sum_{i=1}^N -\log \sqrt{2\pi\sigma^2} - \frac{(z^{(i)} - \mu)^2}{2\sigma^2}$$

$$\ell(\theta; \mathcal{D}) = \sum_i -\log \sqrt{2\pi\sigma^2} - \frac{(y^{(i)} - \theta^T x^{(i)})^2}{2\sigma^2}$$

$$J(\theta; \mathcal{D}) = -\ell(\theta; \mathcal{D}) \propto \sum_{i=1}^N (y^{(i)} - \theta^T x^{(i)})^2$$

Just like MSE!

Additional Slides

Probability Primer

Probability Vocab

Outcomes

Sample space

Events

Probability

Random variable

Discrete random variable

Continuous random variable

Probability mass function

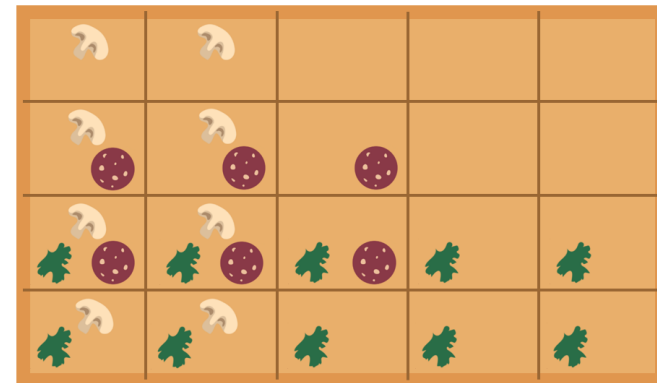
Probability density function

Parameters

Example

$$S \in \{0, 1\}$$
$$S: \Omega \rightarrow \{0, 1\}$$

Random variable for spinach or no spinach



Ω
events

Random variable

R

Distribution

0

1

Sample space

Values

Probabilities or
Densities

Example

Random variable for spinach or no spinach

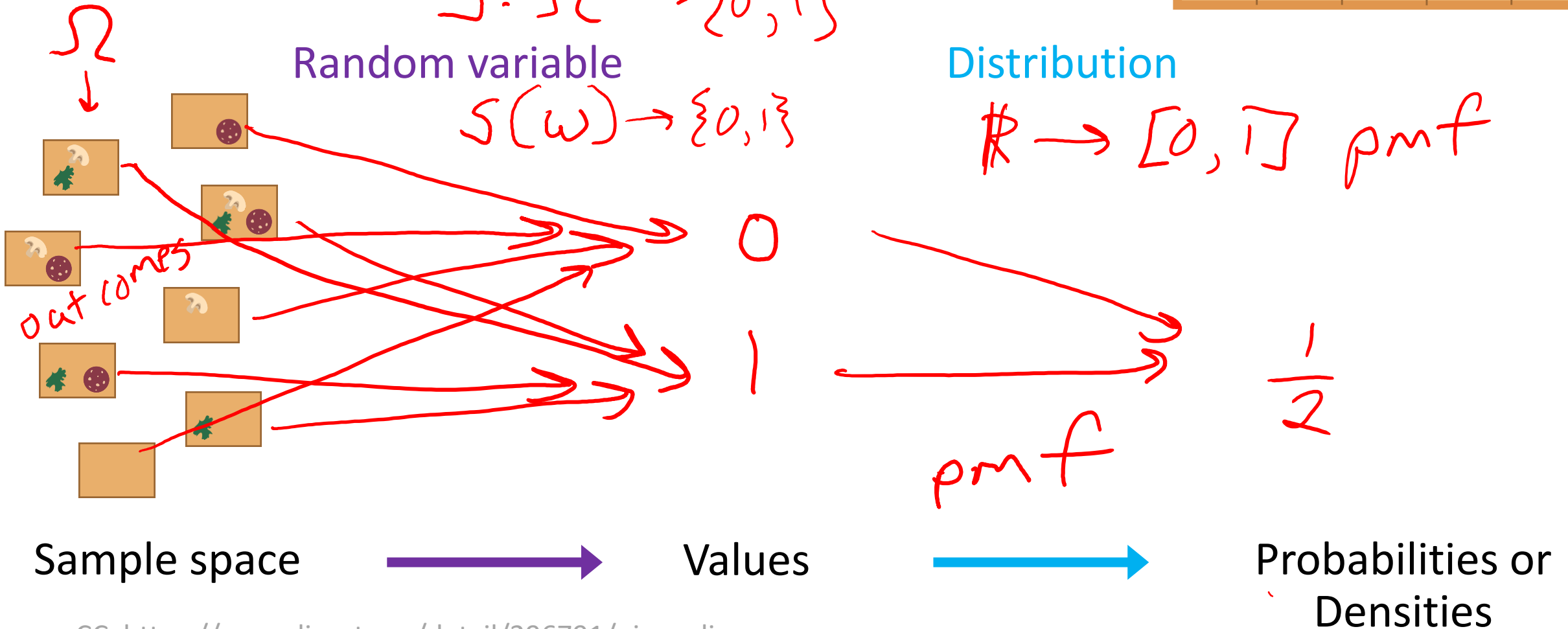
$$S: \Omega \rightarrow \{0, 1\}$$

Random variable

$$S(\omega) \rightarrow \{0, 1\}$$

Distribution

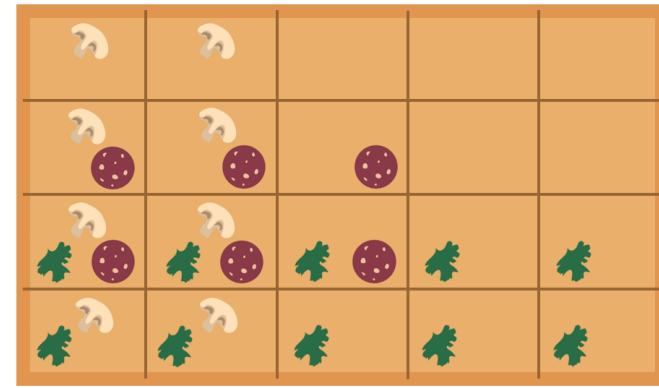
$$\mathbb{R} \rightarrow [0, 1] \text{ pmf}$$



Example

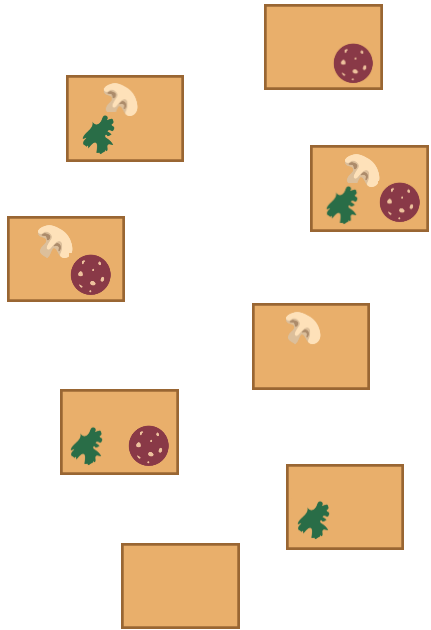
Random variable for topping type with three categories: none, non-meat, meat

time



Random variable

Distribution



Sample space



Values

1
2
3

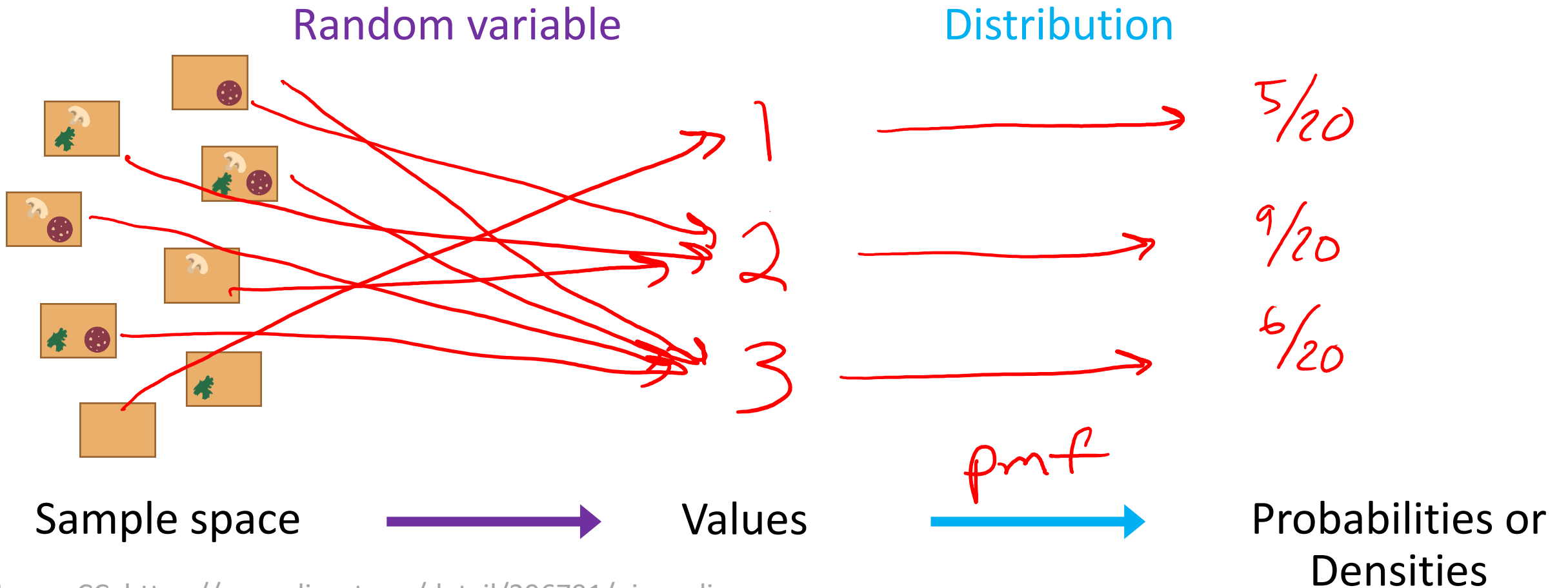
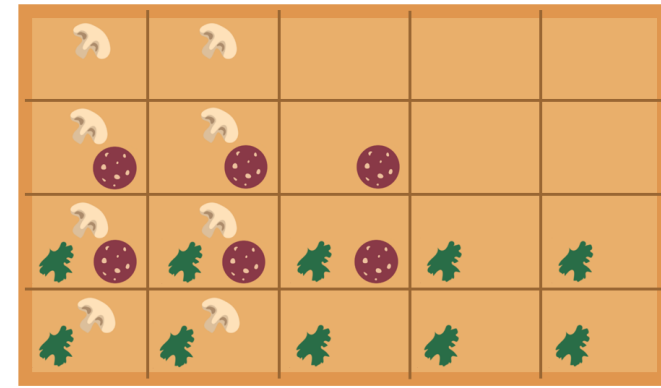
pmf



Probabilities or
Densities

Example

Random variable for topping type with three categories: none, non-meat, meat

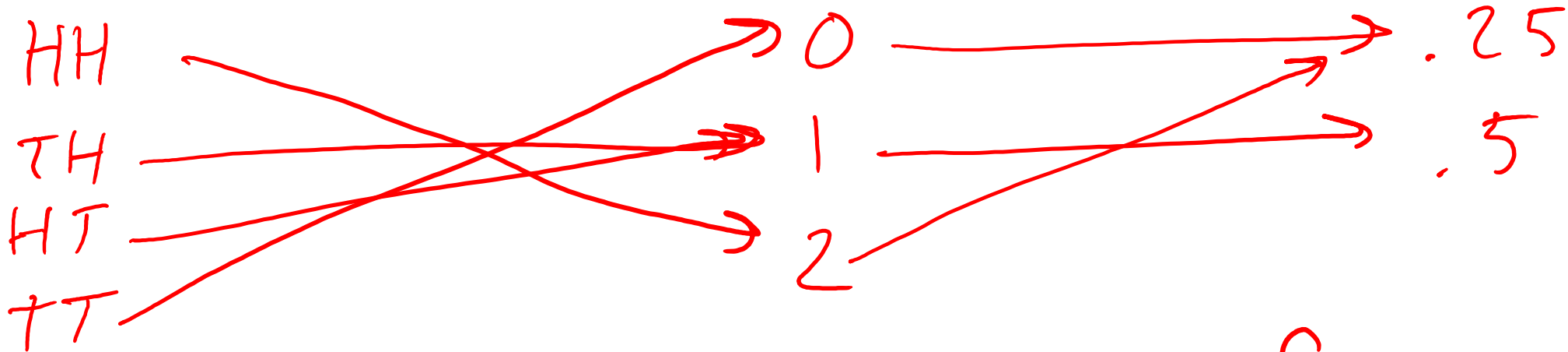


Example

Random variable for number of heads after
two flips of a fair coin

Random variable

Distribution



Sample space



Values



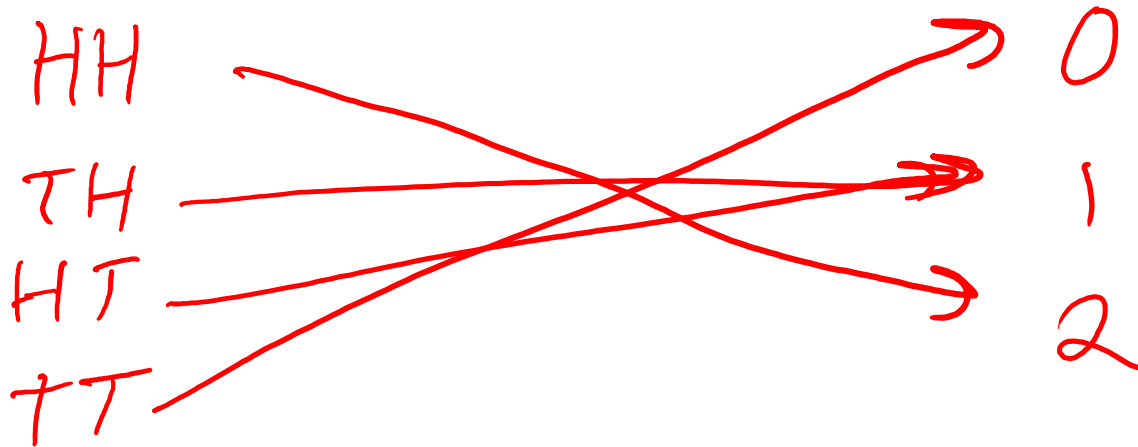
Probabilities or
Densities

Example

Random variable for number of heads after two flips of a *biased* coin that lands heads 75%

Random variable

Distribution



$$\left[\begin{array}{l} \frac{1}{4} \cdot \frac{1}{4} \\ \frac{1}{4} \cdot \frac{3}{4} + \frac{3}{4} \cdot \frac{1}{4} \\ \frac{3}{4} \cdot \frac{3}{4} \end{array} \right]$$

Sample space



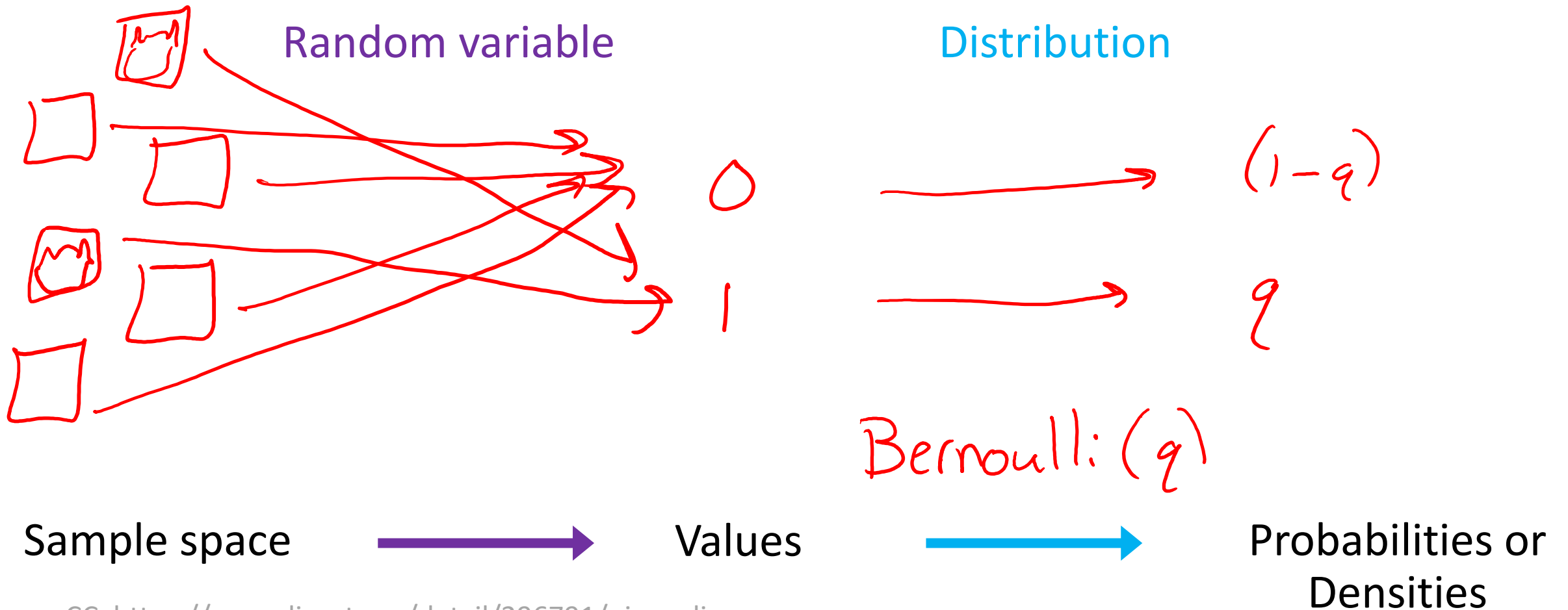
Values



Probabilities or
Densities

Example

Random variable for cat in picture or not



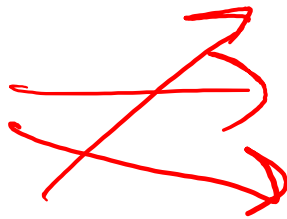
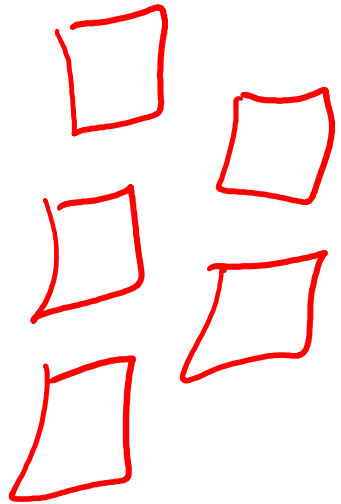
Example

Random variable for animal species in picture assuming one animal picture and available species: dog, cat, pig

Random variable

Distribution

$$\checkmark \sum_{i=1}^3 q_i = 1$$



1
2
3

q_1

q_2

$q_3 = (1 - q_1 - q_2)$

Sample space



Values



Categorical (q_1, q_2, q_3)

Probabilities or
Densities

Example

Random variable for height of student

$$H(\omega) \rightarrow \mathbb{R}$$

Random variable

Continuous

$$f(x) \rightarrow \mathbb{R}^+$$

Distribution

Gaussian
 $N(\mu, \sigma^2)$

$f(x; \mu, \sigma^2)$
pdf
parameters



$\mathbb{R}$

Sample space



Values



Probabilities or
Densities

Probability Vocab

Outcomes

Sample space

Events

Probability

Random variable

Discrete random variable

Continuous random variable

Probability mass function

Probability density function

Parameters

$$\begin{aligned} & \left[\begin{aligned} & \Pr[\text{event}] \\ & P(\text{event}) \end{aligned} \right] \\ & \rightarrow P(\text{value}) \\ & \rightarrow P(Y=1) \end{aligned}$$

$$p(x) = P(X=x)$$

$$f(x)$$

$$p(x)$$

$$\begin{aligned} & \text{CDF} \\ & F(x) = P(X \leq x) \end{aligned}$$

Example Discrete Distributions

Bernoulli

$$P(Y=1) = \phi$$

$$Y \in \{0, 1\} \quad P(Y=0) = 1 - \phi$$

Categorical

$$Y = \begin{bmatrix} Y_1 \\ \vdots \\ Y_K \end{bmatrix}$$

$$Y_k \in \{0, 1\}$$

$$\begin{bmatrix} P(Y_1=1) = \phi_1 \\ P(Y_2=1) = \phi_2 \\ \vdots \\ P(Y_K=1) = \phi_K \end{bmatrix}$$

$$1 = \sum_{k=1}^K \phi_k$$

Binomial

Multinomial

Uniform

Bern Binom.
Categorical Multinom

Example Continuous Distributions

Gaussian

$$p(y; \mu, \sigma^2) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(y-\mu)^2}{2\sigma^2}}$$

Beta

Laplace

Probability Vocab

Marginal

$$p(x) = \sum_y p(x, y) \quad \text{Marginalizing}$$

Joint

$$p(x, y) \quad p(x, y, z, w)$$

Conditional

$$p(x | y) \quad p(x, y | z, w)$$

Notation

Dataset

Parameters, generically θ

$$p(\mathcal{D} \mid \theta), p(\mathcal{D} ; \theta)$$

Random variables

Capital

Values

lower case

$$P(Y=y)$$

A
Σ

Random variable: function that maps events to values

Y is rand variable that maps the event of a coin toss being heads to value one and the event of a coin toss being tails to zero

$$P(Y = 1 \mid \phi) = 3/4, \text{ where } \phi = 3/4$$

$$P(Y = 1) = 3/4$$

Sometimes even

$$P(Y = \text{heads}) = 3/4$$

Probability Toolbox

- Algebra
- Three axioms of probability
- Theorem of total probability
- Definition of conditional probability
- Product rule
- Bayes' theorem
- Chain rule
- Independence
- Conditional independence

Probability Tools Summary

Adding to the toolbox

1. Definition of conditional probability

$$P(A|B) = \frac{P(A, B)}{P(B)}$$

2. Chain Rule *Product rule*

$$\underline{P(A, B)} = P(A | B)P(B)$$

3. Bayes' theorem

$$P(B|A) = \frac{P(A | B)P(B)}{P(A)}$$

4. Chain Rule...

$$P(A_1, \dots, A_N) = P(A_1) \sum_{i=2}^N P(A_i | A_{i-1})$$