

Announcements

- Anonymous feedback form
- Recitation on Friday Sept 10 – Convexity review
- QnA1 due TODAY
- HW1 to be released TODAY

Recap

- **Bayes classifier** – assumes P_{XY} known, optimal for 0/1 loss

$$\begin{aligned} f(X) &= \arg \max_{Y=y} P(Y = y | X = x) \\ &= \arg \max_{Y=y} \underbrace{P(X = x | Y = y)}_{\substack{\text{Class conditional} \\ \text{Distribution of features}}} \underbrace{P(Y = y)}_{\substack{\text{Class distribution} \\ \underline{\underline{}}}} \end{aligned}$$

$P(f(X) \neq Y)$
 $= E_{XY}[1_{f(X) \neq Y}]$

- **Gaussian Bayes classifier** – assumes
Class distribution is Bernoulli/Multinomial
Class conditional distribution of features is Gaussian
- **Decision boundary** – (binary classification)

d-dim Gaussian Bayes classifier

$$f(X) = \arg \max_{Y=y} \underbrace{P(X = x|Y = y)}_{\text{Class conditional}} \underbrace{P(Y = y)}_{\text{Class distribution}}$$

➤ What decision boundaries can we get in d-dim?

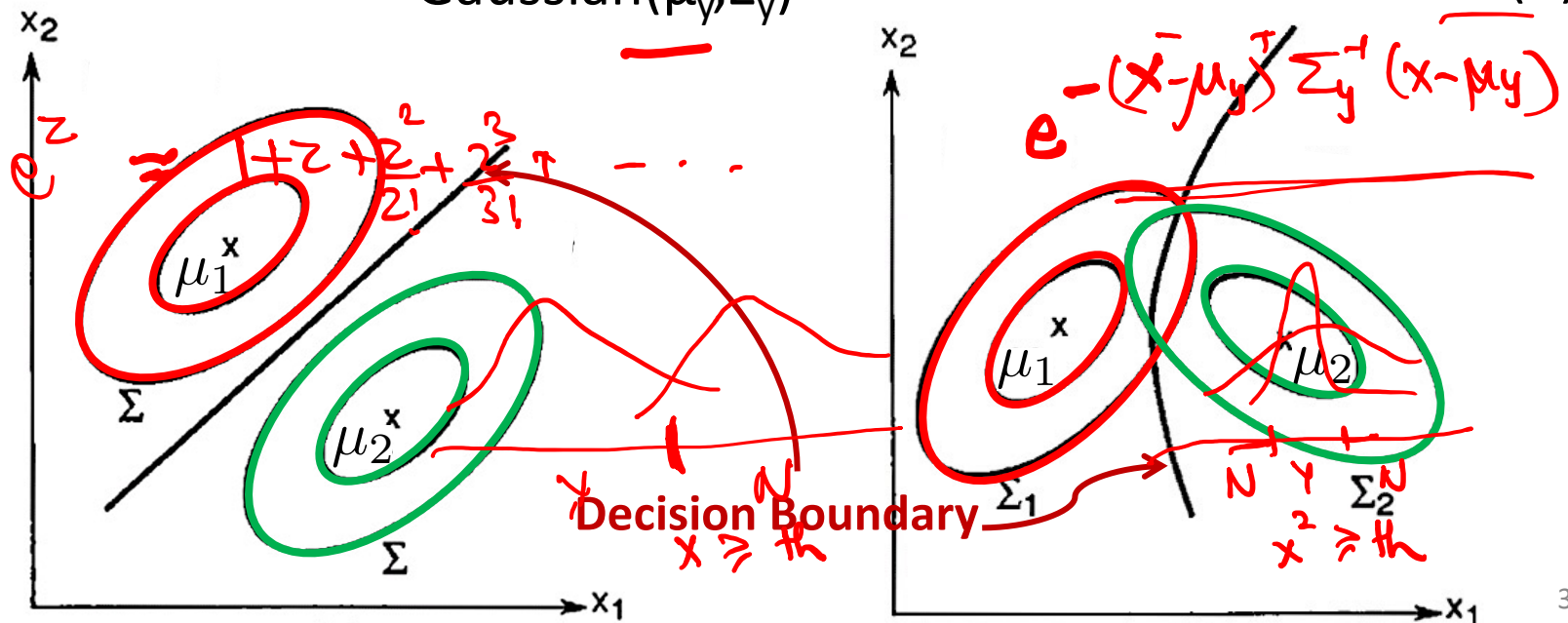
Class conditional

Class distribution

Distribution of features

Gaussian(μ_y, Σ_y)

Bernoulli(θ)



Decision Boundary of Gaussian Bayes

- Decision boundary is set of points x : $P(Y=1|X=x) = P(Y=0|X=x)$

Compute the ratio

$$\begin{aligned}
 1 &= \frac{P(Y=1|X=x)}{P(Y=0|X=x)} = \frac{\overset{\sim N(\mu_1, \Sigma_1)}{P(X=x|Y=1)} \overset{\theta}{P(Y=1)}}{\underset{\sim N(\mu_0, \Sigma_0)}{P(X=x|Y=0)} \underset{1-\theta}{P(Y=0)}} \quad \text{where } Y \sim \text{Ber}(\theta) \\
 &= \frac{1}{\sqrt{2\pi|\Sigma_1|}} \frac{e^{-\frac{1}{2}(x-\mu_1)^T \Sigma_1^{-1} (x-\mu_1)}}{e^{-\frac{1}{2}(x-\mu_0)^T \Sigma_0^{-1} (x-\mu_0)}} \frac{\sqrt{2\pi|\Sigma_0|}}{1-\theta} \frac{\theta}{1-\theta} \\
 &= \sqrt{\frac{|\Sigma_0|}{|\Sigma_1|}} \exp \left(-\frac{(x-\mu_1)^T \Sigma_1^{-1} (x-\mu_1)}{2} + \frac{(x-\mu_0)^T \Sigma_0^{-1} (x-\mu_0)}{2} \right) \frac{\theta}{1-\theta}
 \end{aligned}$$

Handwritten notes: $\frac{e^a}{e^b} = e^{a-b}$, quadratic in x, $-x^T \Sigma_1^{-1} x + x^T \Sigma_0^{-1} x$

In general, this implies a quadratic equation in x . But if $\Sigma_1 = \Sigma_0$, then quadratic part cancels out and decision boundary is linear.

Glossary of Machine Learning

- Feature/Attribute
- iid
- Bayes classifier
- Class distribution
- Class conditional distribution of features
- Decision boundary

Maximum Likelihood Estimation (MLE)

Aarti Singh

Machine Learning 10-315

Sept 8, 2021



MACHINE LEARNING DEPARTMENT



How to learn parameters from data?

MLE

(Discrete case)

Learning parameters in distributions

$$P(Y = \bullet) = \theta$$

$$P(Y = \bullet) = 1 - \theta$$

Learning θ is equivalent to learning probability of head in coin flip.

➤ How do you learn that?

Data =



Answer: 3/5

➤ Why??

Bernoulli distribution

Data, $D =$



- Parameter θ : $P(\text{Heads}) = \theta$, $P(\text{Tails}) = 1-\theta$
- Flips are **i.i.d.**:
 - **Independent** events
 - **Identically distributed** according to Bernoulli distribution

Choose θ that maximizes the probability of observed data
aka Likelihood

Maximum Likelihood Estimation (MLE)

Choose θ that maximizes the probability of observed data (aka likelihood)

$$D = \overset{H}{x_1} \dots \overset{T}{x_n} \quad x_i = \begin{cases} 1 & H \\ 0 & T \end{cases}$$

$$\hat{\theta}_{MLE} = \arg \max_{\theta} P(D; \theta)$$

$$= P(x_1 \dots x_n; \theta)$$

$$= \prod_{i=1}^n P(x_i; \theta)$$

$$= \prod_{i=1}^n \theta^{x_i} (1-\theta)^{1-x_i} = \theta^{\sum x_i} (1-\theta)^{n - \sum x_i}$$

$\max_{\theta}$

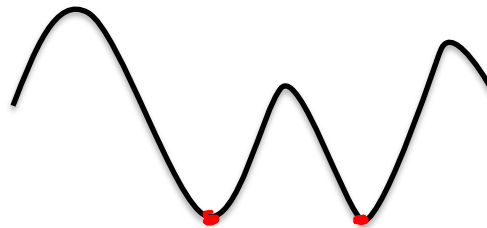
MLE of probability of head:

$$\hat{\theta}_{MLE} = \frac{\alpha_H}{\alpha_H + \alpha_T} = 3/5$$

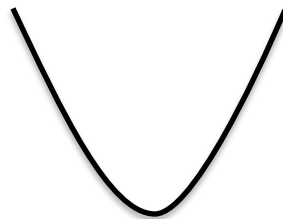
"Frequency of heads"

Short detour - Optimization

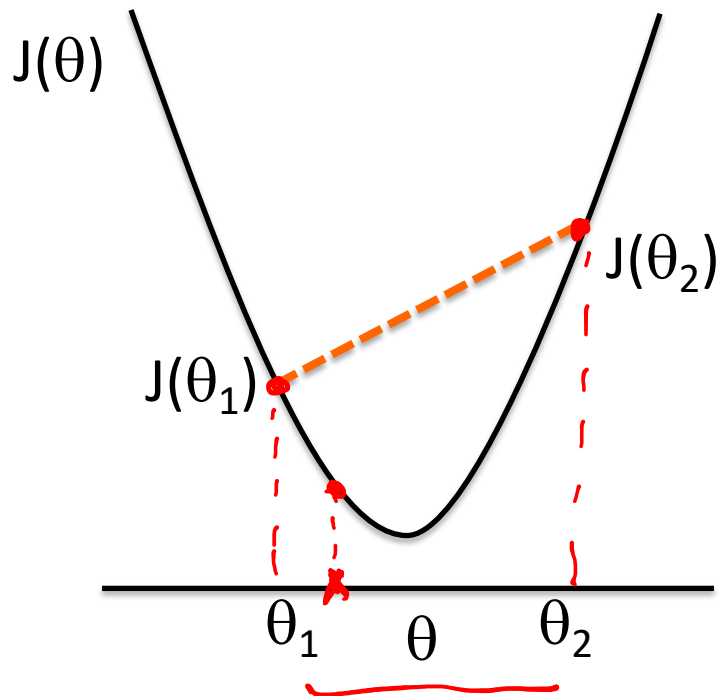
- Optimization objective $J(\theta)$
- Minimum value $J^* = \min_{\theta} J(\theta)$
- Minima (points at which minimum value is achieved) may not be unique



- If function is strictly convex, then minimum is unique

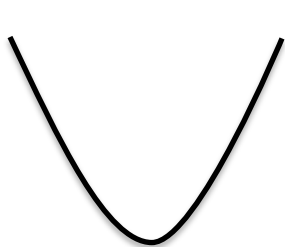


Convex functions

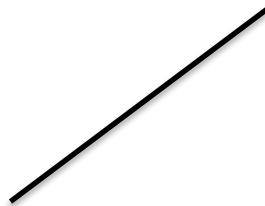


A function $J(\theta)$ is called **convex** if the line joining two points $J(\theta_1), J(\theta_2)$ on the function does not go below the function on the interval $[\theta_1, \theta_2]$

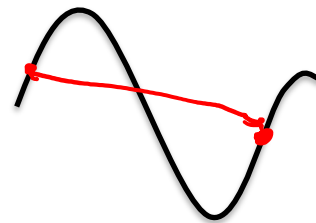
(Strictly) Convex functions have a unique minimum!



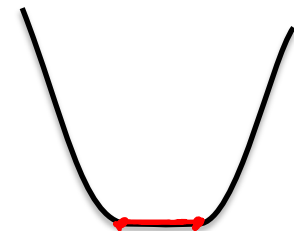
Convex



Both Concave
& Convex



Neither



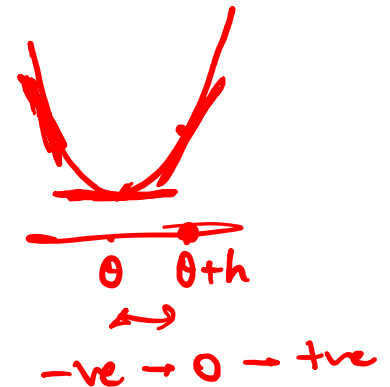
Convex but not
strictly convex

Optimizing convex (concave) functions

- Derivative of a function

$$\frac{\partial J(\theta)}{\partial \theta} = \lim_{h \rightarrow 0} \frac{J(\theta+h) - J(\theta)}{h}$$

– Partial derivative



- Derivative is zero at minimum of a convex function
- Second derivative is positive at minimum of a convex function

Optimizing convex (concave) functions

➤ What about

concave functions?

non-convex/non-concave functions?

functions that are not differentiable?

optimizing a function over a bounded domain aka
constrained optimization?

Bernoulli MLE Derivation

$$\hat{\theta}_{MLE} = \arg \max_{\theta} \overline{P(D | \theta)}$$

$$= \arg \max_{\theta} \prod_{i=1}^n \theta^{x_i} (1-\theta)^{1-x_i}$$

$$= \arg \max_{\theta} \underbrace{\sum_{i=1}^n x_i \log \theta + (1-x_i) \log(1-\theta)}_{J(\theta)}$$

$$\rightarrow \frac{\partial J(\theta)}{\partial \theta} = \sum_{i=1}^n \frac{x_i}{\theta} - \frac{(1-x_i)}{1-\theta} = 0$$

$$x_i = \begin{cases} 1 & H \\ 0 & T \end{cases}$$

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

$$\frac{\alpha_H}{\theta} = \frac{\alpha_T}{1-\theta}$$

$$\Rightarrow \alpha_H(1-\theta) = \alpha_T \theta$$

$$\alpha_H = (\alpha_H + \alpha_T) \theta =$$

$$\theta = \frac{\alpha_H}{\alpha_H + \alpha_T}$$

Multinomial distribution

Data, D = rolls of a dice



- $P(1) = p_1, P(2) = p_2, \dots, P(6) = p_6$ $p_1 + \dots + p_6 = 1$
- Rolls are **i.i.d.**:
 - **Independent** events
 - **Identically distributed** according to Multinomial(θ) distribution where

$$\theta = \{p_1, p_2, \dots, p_6\}$$

Choose θ that maximizes the probability of observed data
aka “Likelihood”

Maximum Likelihood Estimation (MLE)

Choose θ that maximizes the probability of observed data

$$\hat{\theta}_{MLE} = \arg \max_{\theta} P(D | \theta)$$

MLE of probability of rolls:

$$\hat{\theta}_{MLE} = \hat{p}_{1,MLE}, \dots, \hat{p}_{6,MLE}$$

$$\hat{p}_{y,MLE} = \frac{\alpha_y \leftarrow \text{Rolls that turn up } y}{\sum_y \alpha_y \leftarrow \text{Total number of rolls}}$$

"Frequency of roll y "

How to learn parameters from data?

MLE

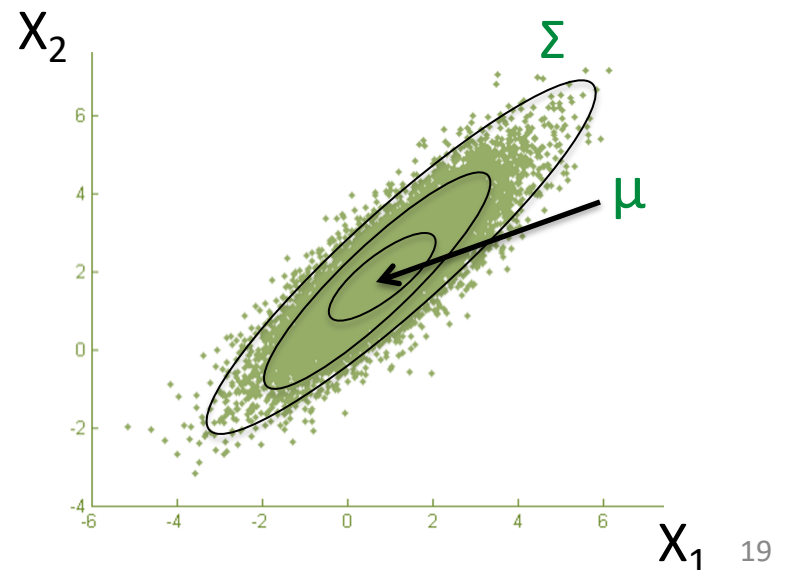
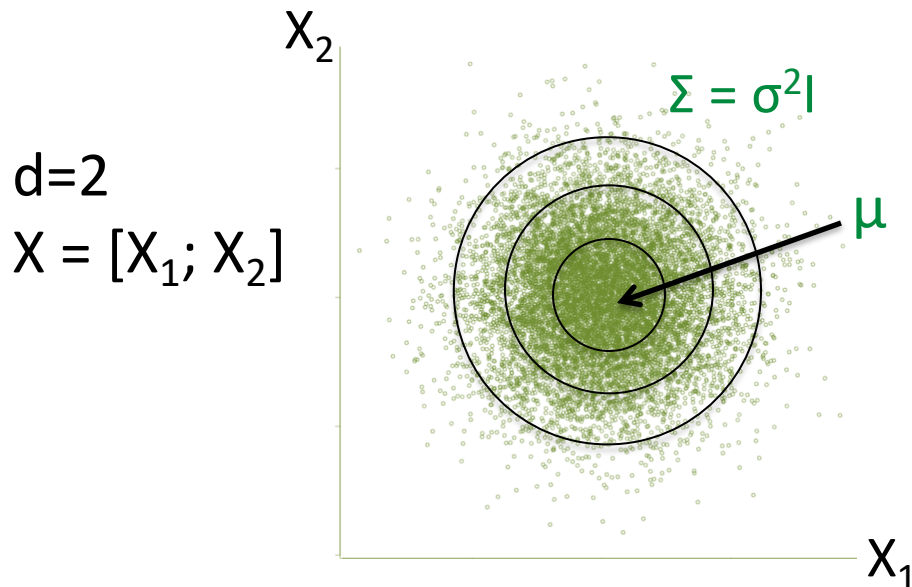
(Continuous case)

d-dim Gaussian distribution

X is Gaussian $N(\mu, \Sigma)$

μ is d-dim vector, Σ is dxd dim matrix

$$P(X = x | \mu, \Sigma) = \frac{1}{\sqrt{(2\pi)^d |\Sigma|}} \exp \left(-\frac{1}{2} (\mathbf{x} - \mu)^T \Sigma^{-1} (\mathbf{x} - \mu) \right),$$

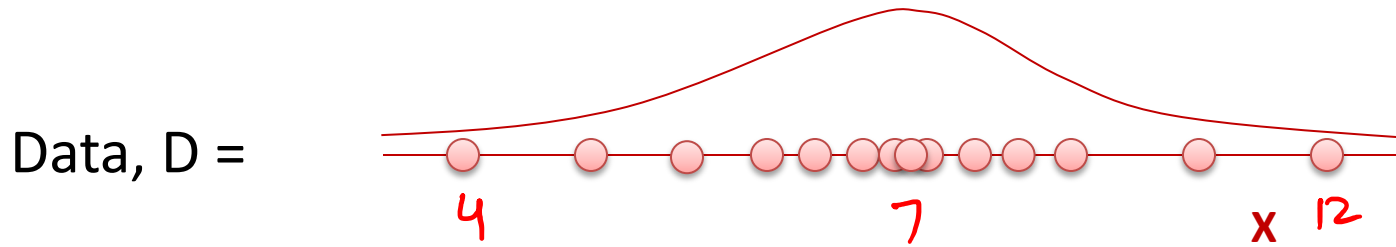


How to learn parameters from data?

MLE

(Continuous case)

Gaussian distribution



- Parameters: μ – mean, σ^2 – variance
- Data are **i.i.d.**:
 - **Independent** events
 - **Identically distributed** according to Gaussian distribution

Maximum Likelihood Estimation (MLE)

Choose $\theta = (\mu, \sigma^2)$ that maximizes the probability of observed data

$$\begin{aligned}\hat{\theta}_{MLE} &= \arg \max_{\theta} P(\overset{\checkmark}{D} | \theta) \\ &= \arg \max_{\theta} \prod_{i=1}^n P(X_i | \theta) \quad \text{Independent draws} \\ &= \underbrace{\mu, \sigma^2}_{\rightarrow} \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(x_i - \mu)^2}{2\sigma^2}}\end{aligned}$$

Maximum Likelihood Estimation (MLE)

Choose $\theta = (\mu, \sigma^2)$ that maximizes the probability of observed data

$$\hat{\theta}_{MLE} = \arg \max_{\theta} P(D | \theta)$$

$$= \arg \max_{\theta} \prod_{i=1}^n P(X_i | \theta) \quad \text{Independent draws}$$

$$= \arg \max_{\theta} \prod_{i=1}^n \frac{1}{\sqrt{2\pi\sigma^2}} e^{-(X_i - \mu)^2 / 2\sigma^2} \quad \text{Identically distributed}$$

Maximum Likelihood Estimation (MLE)

Choose $\theta = (\mu, \sigma^2)$ that maximizes the probability of observed data

$$\hat{\theta}_{MLE} = \arg \max_{\theta} P(D | \theta)$$

$$= \arg \max_{\theta} \prod_{i=1}^n P(X_i | \theta) \quad \text{Independent draws}$$

$$= \arg \max_{\theta} \prod_{i=1}^n \frac{1}{\sqrt{2\pi\sigma^2}} e^{-(X_i - \mu)^2 / 2\sigma^2} \quad \text{Identically distributed}$$

$$= \arg \max_{\theta = (\mu, \sigma^2)} \underbrace{\frac{1}{(2\pi\sigma^2)^{n/2}} e^{-\sum_{i=1}^n (X_i - \mu)^2 / 2\sigma^2}}_{J(\theta)}$$

$$e^a e^b = e^{a+b}$$

MLE for Gaussian mean

➤ Poll

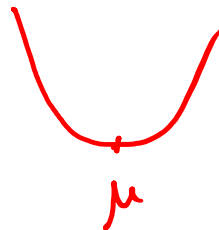
$$P(D|\theta) = \frac{1}{(2\pi\sigma^2)^{n/2}} e^{-\sum_{i=1}^n (X_i - \mu)^2 / 2\sigma^2}$$

A. $\max_{\mu} \sum_{i=1}^n (X_i - \mu)^2$

C. $\max_{\mu} \mu^2 - 2\mu \sum_{i=1}^n X_i$

B. $\min_{\mu} \sum_{i=1}^n (X_i - \mu)^2$

D. $\max_{\mu} n\mu^2 - 2\mu \sum_{i=1}^n X_i$



$$-\sum_{i=1}^n 2(X_i - \mu) = 0$$

$$n\mu = \sum_{i=1}^n X_i$$
$$\mu = \frac{1}{n} \sum_{i=1}^n X_i$$

MLE for Gaussian mean and variance

$$\hat{\mu}_{MLE} = \frac{1}{n} \sum_{i=1}^n x_i \quad \checkmark$$

$$\hat{\sigma}_{MLE}^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \hat{\mu})^2 \quad \checkmark$$

Self exercise:

Derive MLE of variance?

d-dimensional versions?

$\hat{\mu} = \begin{bmatrix} \hat{\mu}_1 \\ \hat{\mu}_2 \\ \vdots \\ \hat{\mu}_d \end{bmatrix} = \begin{bmatrix} \frac{1}{n} \sum_{i=1}^n x_{i1} \\ \vdots \\ \frac{1}{n} \sum_{i=1}^n x_{id} \end{bmatrix}$
MLE for uniform or exponential distribution?



More coming up in HW1

Unbiased estimates

- Bias of an estimate $E[\hat{\theta}] - \theta$
- Unbiased estimate if $E[\hat{\theta}] = \theta$

$$\hat{\mu}_{MLE} = \frac{1}{n} \sum_{i=1}^n X_i$$

$$E[\hat{\mu}_{MLE}] = \frac{1}{n} \sum_{i=1}^n E[X_i] \\ \quad \quad \quad \mu \\ = \mu$$

➤ Is the MLE of mean unbiased?

➤ Is the MLE of variance unbiased?

$$E[\hat{\sigma}_{MLE}^2] = \sigma^2 \frac{n-1}{n}$$

➤ How can you make it unbiased?