



Artificial Intelligence III:  
Artificial Intelligence and Deep Learning

*Lecture 3*

# Regression & Classification

Dr. Patrick Chan

[patrickchan@ieee.org](mailto:patrickchan@ieee.org)

South China University of Technology, China



## Agenda

- ◆ Supervised Learning
  - Regression
    - Least Mean Square
  - Classification
    - Probability



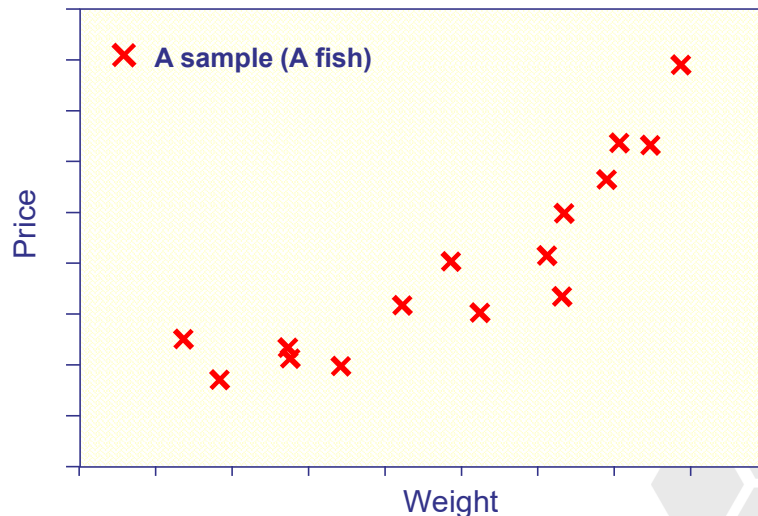


# Regression

- ◆ Given 15 fishes: weight and prices
- ◆ Objective: Predict the price of a fish

| Weight (kg) | Price (\$) |
|-------------|------------|
| 2.2         | 20         |
| 2.6         | 31         |
| 1.2         | 16.5       |
| 0.7         | 10         |

⋮



3

Dr. Patrick Chan @ SCUT

Artificial Intelligence III: Artificial Intelligence and Deep Learning - Lecture 3



# Regression

- ◆ The  $i^{\text{th}}$  training sample:  $(x^{(i)}, y^{(i)})$ 
  - $x^{(i)} = [x_1^{(i)}, x_2^{(i)}, \dots, x_d^{(i)}]^T \in X$  : **Feature vector**
    - $x_j^{(i)}$  : the  $j^{\text{th}}$  feature
    - $X$  : the input space
  - $y^{(i)} \in Y$  : **Target Value**
    - $Y$  : the output space

|                      | $x$         | $y$        |
|----------------------|-------------|------------|
|                      | Weight (kg) | Price (\$) |
| $(x^{(1)}, y^{(1)})$ | 2.2         | 20         |
| $(x^{(2)}, y^{(2)})$ | 2.6         | 31         |
| $(x^{(3)}, y^{(3)})$ | 1.2         | 16.5       |
| $(x^{(4)}, y^{(4)})$ | 0.7         | 10         |
|                      | ⋮           |            |

- ◆ **Training set:**  $\{(x^{(i)}, y^{(i)}) \mid i = 1 \dots n\}$ 
  - $n$  is the number of training samples

4

Dr. Patrick Chan @ SCUT

Artificial Intelligence III: Artificial Intelligence and Deep Learning - Lecture 3



# Regression

- ◆ Train a function  $h_{\theta}(x)$  to predict  $y$ 
  - $\theta$  is the parameter vectors of the model
    - Different parameters yields different predictions (different  $h$ )
    - Example: weight
  - $h_{\theta}: X \rightarrow Y$ , mapping from  $X$  to  $Y$
  - $h_{\theta}$  is called a predictor or hypothesis
- ◆ Objective: Build a “good”  $h_{\theta}$ 
  - What does “good” mean?



## Regression

# Objective Function

- ◆ Objective: the predicted value on a training sample closer to the real one
  - Smaller difference between  $h_{\theta}(x^{(i)})$  and  $y^{(i)}$
- ◆ Cost function (objective function)
  - May contains other terms besides Error
  - Mean Square Error is a classic measure

$$J(\theta) = \frac{1}{2n} \sum_{i=1}^n \underbrace{(h_{\theta}(x^{(i)}) - y^{(i)})^2}_{\text{mean square error}}$$

square error

- Error is a distance measure
- Square avoids the cancellation of positive and negative error





# LMS Algorithm

- ◆ **Least Mean Squares (LMS)** aims to minimize  $J(\theta)$  by adjusting  $\theta$

$$J(\theta) = \frac{1}{2n} \sum_{i=1}^n (h_{\theta}(x^{(i)}) - y^{(i)})^2$$

- ◆ ML is closely related to Optimization problem
  - Usually, the optimization is quite complicated
    - Since each parameter is a variable
    - Iterative method is used

7

Dr. Patrick Chan @ SCUT

Artificial Intelligence III: Artificial Intelligence and Deep Learning - Lecture 3



## LMS Algorithm

# Gradient Descent

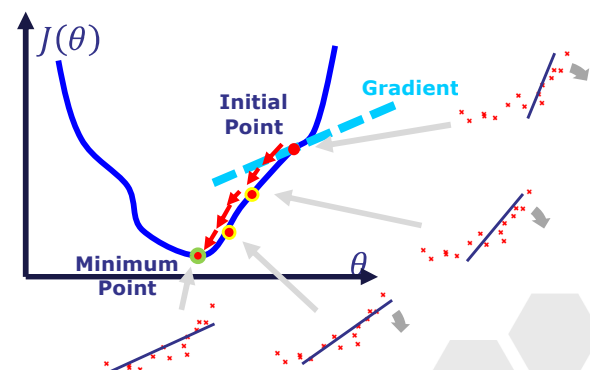
- ◆ When  $h_{\theta}$  is **differentiable**, **gradient descent** can be used to minimize  $J(\theta)$

$$J(\theta) = \frac{1}{2n} \sum_{i=1}^n (h_{\theta}(x^{(i)}) - y^{(i)})^2$$

- ◆ Influence on  $J(\theta)$  by changing the parameter ( $\theta$ ) slightly

$$\theta^{(t+1)} = \theta^{(t)} - \alpha \frac{\partial J(\theta^{(t)})}{\partial \theta}$$

- $\alpha$  : the learning rate
- $\theta^{(t)}$  : the parameter at the time  $t$



8

Dr. Patrick Chan @ SCUT

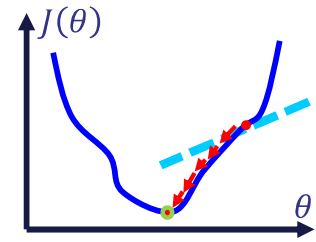
Artificial Intelligence III: Artificial Intelligence and Deep Learning - Lecture 3





## ◆ Algorithm

- Start with an arbitrarily chosen weight  $\theta^{(1)}$
- Let  $t = 0$
- Loop
  - $t = t + 1$
  - Compute gradient vector  $\partial J(\theta^{(t)})/\partial \theta$
  - Next value  $\theta^{(t+1)}$  determined by moving some distance from  $\theta^{(t)}$  in the direction of the steepest descent



$$\theta^{(t+1)} = \theta^{(t)} - \alpha \frac{\partial}{\partial \theta} J(\theta^{(t)})$$

- i.e., along the negative of the gradient
- Until Finish Training



- ◆ Recall,  $\theta = [\theta_1, \theta_2, \dots, \theta_m]$
- ◆ Updated Rule for the  $j^{\text{th}}$  parameter

$$\theta_j^{(t+1)} = \theta_j^{(t)} - \alpha \frac{\partial}{\partial \theta_j} J(\theta_j^{(t)})$$

- ◆ All parameters should be updated at the same time





- How to calculate  $\frac{\partial}{\partial \theta_j} J(\theta)$ ?

$$\theta_j = \theta_j - \alpha \frac{\partial}{\partial \theta_j} J(\theta)$$

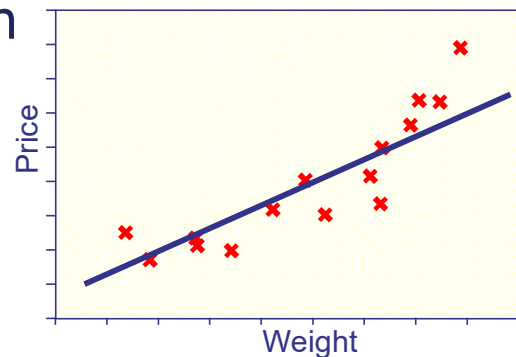
$$\begin{aligned} J(\theta) &= \frac{1}{2n} \sum_{i=1}^n (h_{\theta}(x^{(i)}) - y^{(i)})^2 \\ \frac{\partial}{\partial \theta_j} J(\theta) &= \frac{\partial}{\partial \theta_j} \frac{1}{2n} \sum_{i=1}^n (h_{\theta}(x^{(i)}) - y^{(i)})^2 \\ &= \frac{1}{2n} \sum_{i=1}^n \frac{\partial}{\partial \theta_j} (h_{\theta}(x^{(i)}) - y^{(i)})^2 \\ &= \frac{1}{2n} \sum_{i=1}^n 2(h_{\theta}(x^{(i)}) - y^{(i)}) \frac{\partial}{\partial \theta_j} (h_{\theta}(x^{(i)}) - y^{(i)}) \\ &= \frac{1}{2n} \sum_{i=1}^n 2(h_{\theta}(x^{(i)}) - y^{(i)}) \left( \frac{\partial h_{\theta}(x^{(i)})}{\partial \theta_j} \right) \end{aligned}$$

Depend on a model



- Example: Linear Function

$$\begin{aligned} h_{\theta}(x) &= \theta_1 x_1 + \theta_2 x_2 + \cdots + \theta_d x_d \\ &= \sum_{i=1}^d \theta_i x_i \end{aligned}$$

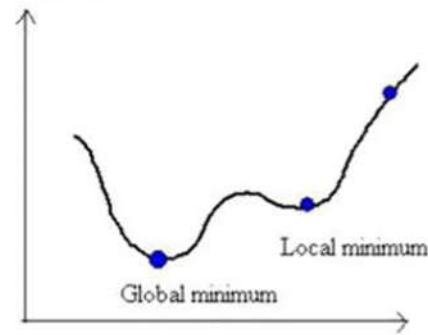


$$\begin{aligned} \frac{\partial}{\partial \theta_j} J(\theta) &= \frac{1}{2n} \sum_{i=1}^n 2(h_{\theta}(x^{(i)}) - y^{(i)}) \frac{\partial h_{\theta}(x^{(i)})}{\partial \theta_j} \\ &= \frac{1}{2n} \sum_{i=1}^n 2(h_{\theta}(x^{(i)}) - y^{(i)}) \frac{\partial}{\partial \theta_j} \left( \sum_{k=1}^d \theta_k x_k^{(i)} \right) \\ &= \frac{1}{n} \sum_{i=1}^n (h_{\theta}(x^{(i)}) - y^{(i)}) x_j^{(i)} \end{aligned}$$



## ◆ Related Issues:

- Size of Learning Rate ( $\alpha$ )
  - Too **small**, convergence is needlessly **slow**
  - Too **large**, the correction process will **overshoot** and **cannot even diverge**
- Sub-optimal Solution
  - **Trapped by local minimum**
- We will study Gradient Descent again in Artificial Neural Network



# Classification

- ◆ Objective: output a class based on the features of a sample

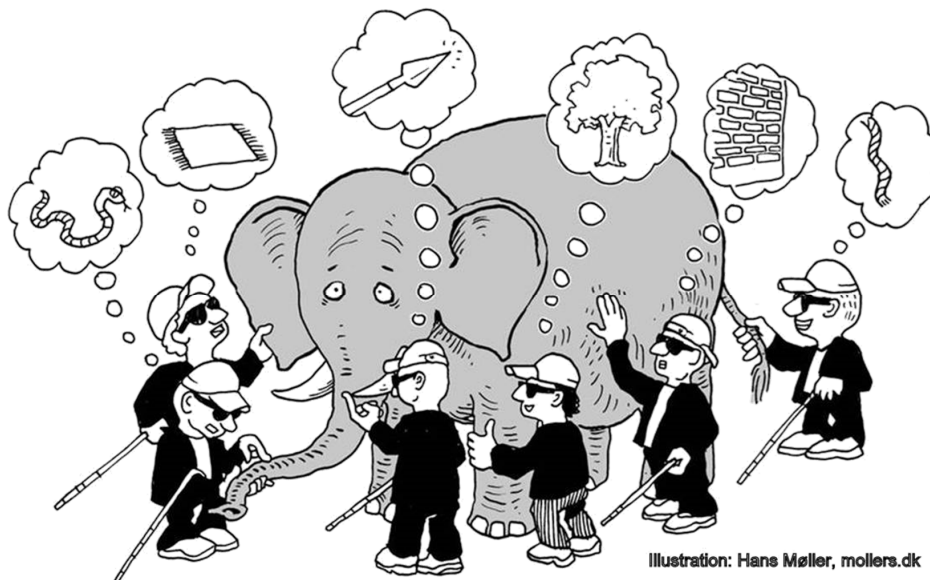


Illustration: Hans Møller, mollers.dk



- ◆ Peter went to body check to see if he is ok

$y = (\text{ill}, \text{healthy})$

- ◆ According to the previous records, the doctor concluded

- 85% of people was healthy

$$P(y = \text{healthy}) = 0.85$$

- 15% of people was ill

$$P(y = \text{ill}) = 0.15$$

- Therefore, Peter was healthy

$$P(y = \text{healthy}) > P(y = \text{ill})$$

$y$

| Person | Status  |
|--------|---------|
| A      | Ill     |
| B      | Healthy |
| C      | Healthy |
| D      | Ill     |
| ⋮      |         |

- ◆ Should Peter be satisfied with this diagnosis?

- This decision is based on **Prior Probability**  $P(y)$



- ◆ Physical condition of persons should be considered

- Quantify the characteristics (features), denoted by  $x$
- E.g. red blood cell #, white blood cell #, temperature

- ◆ Assume only "white blood cell #" is measured

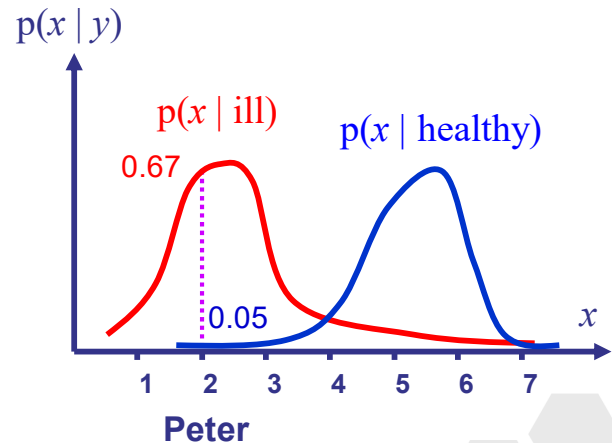
$x$                        $y$

| Person | White Blood Cell # | Status ( $y$ ) |
|--------|--------------------|----------------|
| A      | 50                 | Ill            |
| B      | 42                 | Healthy        |
| C      | 39                 | Healthy        |
| D      | 62                 | Ill            |
| ⋮      |                    |                |



## Classification Likelihood

- ◆ Assume the white blood cell # ( $x$ ) of Peter is 2
- ◆ A probability density function (pdf) of persons is considered
- ◆ The Doctor said
  - $p(x=2 \mid \text{ill}) = 0.67$
  - $p(x=2 \mid \text{healthy}) = 0.05$
  - Therefore, Peter is ill
- ◆ Should we be satisfied?
  - This decision is based on **Likelihood**  $p(x|y)$



## Classification Posterior Probability

- ◆ Using Prior Probability ( $P(y)$ ) or Likelihood ( $p(x|y)$ ) is not suitable
- ◆ **Posterior Probability** is a **better** choice  
 $P(y|x)$  : given  $x$ , the probability of  $y$
- ◆ **Bayes Decision Rule (Bayes Classifier)**
  - When  $P(y_1|x) > P(y_2|x)$ ,  $x$  is  $y_1$
  - When  $P(y_2|x) > P(y_1|x)$ ,  $x$  is  $y_2$
  - When  $P(y_1|x) = P(y_2|x)$ , no decision
- ◆ **How to obtain  $P(y|x)$ ?**
  - Obtaining from data is difficult as  $x$  is usually a continuous value



## ◆ Bayes Formula

$$\overset{\text{posterior}}{P(y|x)} = \frac{\overset{\text{likelihood}}{p(x|y)} \overset{\text{prior}}{P(y)}}{\underset{\text{evidence}}{p(x)}}$$

- ◆ Likelihood and prior probability may be estimated by using a dataset (*Discuss it later in the lecture*)
- ◆ How about evidence  $p(x)$ ?



- ◆  $p(x)$  is difficult to be obtained relatively
  - $p(x)$  contain all kinds of samples, which is more complicated than  $p(x|y)$
  - It can be neglected in decision making

- ◆  $x$  is classified as  $y_1$  if

$$p(y_1|x) > p(y_2|x)$$

$$\frac{p(x|y_1)P(y_1)}{p(x)} > \frac{p(x|y_2)P(y_2)}{p(x)}$$

$$p(x|y_1)P(y_1) > p(x|y_2)P(y_2)$$

$$P(y|x) = \frac{p(x|y)P(y)}{p(x)}$$

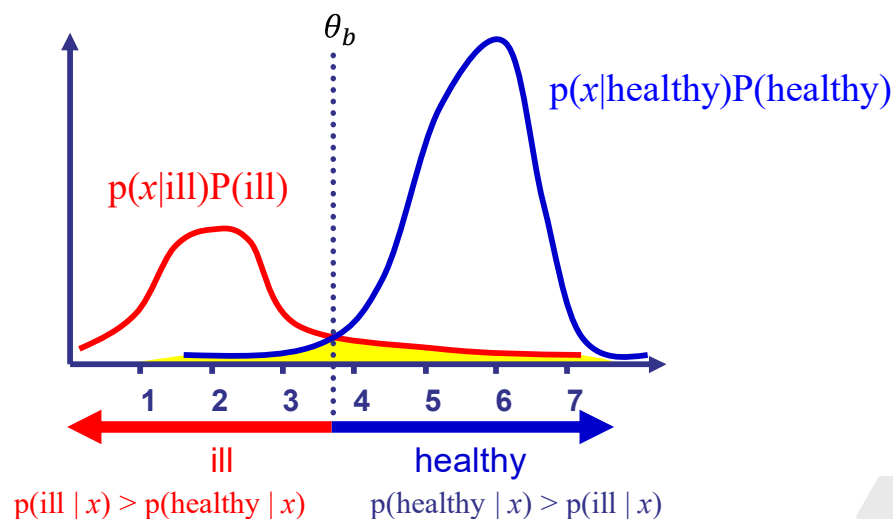




- ◆ Given:  $p(x=2 | \text{ill}) = 0.67$        $p(x=2 | \text{healthy}) = 0.05$   
 $P(\text{ill}) = 0.15$        $P(\text{healthy}) = 0.85$
- ◆ Recall, Bayes Decision Rule
  - Decide  $y_1$  if  $P(y_1 | x) > P(y_2 | x)$
  - Decide  $y_2$  if  $P(y_2 | x) > P(y_1 | x)$
- ◆  $P(\text{healthy} | x = 2) \propto p(x=2 | \text{healthy}) \times P(\text{healthy})$   
 $= 0.05 \times 0.85 = 0.0425$
- ◆  $P(\text{ill} | x = 2) \propto p(x=2 | \text{ill}) \times P(\text{ill})$       \* Note that if  $p(x)$  is considered, then  $P(y_1|x) + P(y_2|x) = 1$ .  
 $= 0.67 \times 0.15 = 0.1005$
- ◆  $0.1005 > 0.0425$ , therefore, **Peter is ill**



- ◆ Recall, Bayes Decision Rule:
  - if  $P(y_1 | x) > P(y_2 | x)$ , decide  $y_1$ ; otherwise decide  $y_2$
- ◆ Its Decision Boundary:





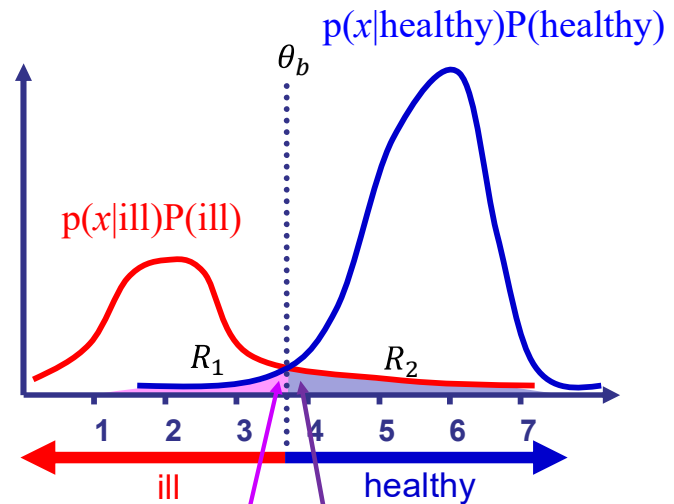
# Classification: Bayes Decision Rule

## Decision Boundary

### ◆ Error of Bayes Decision Rule

|      |       | Prediction                    |                               |
|------|-------|-------------------------------|-------------------------------|
|      |       | $\hat{y}_1$                   | $\hat{y}_2$                   |
| True | $y_1$ | Predict: $y_1$<br>True: $y_1$ | Predict: $y_2$<br>True: $y_1$ |
|      | $y_2$ | Predict: $y_1$<br>True: $y_2$ | Predict: $y_2$<br>True: $y_2$ |

■ Correct  
■ Wrong



Samples is healthy  
when prediction is ill

Samples is ill  
when prediction is healthy

23

Dr. Patrick Chan @ SCUT

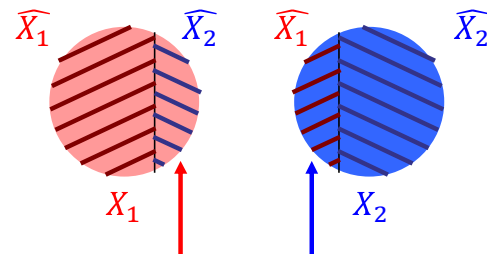
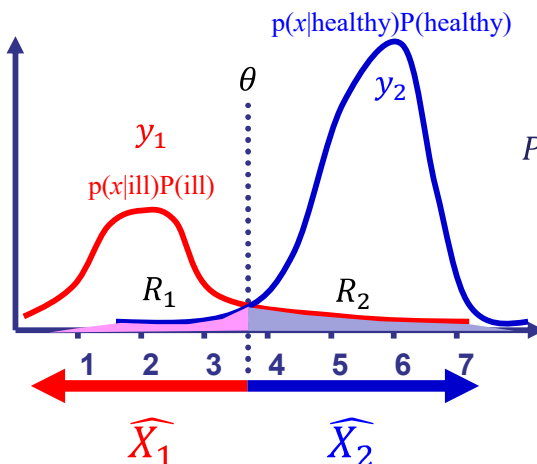
Artificial Intelligence III: Artificial Intelligence and Deep Learning - Lecture 3



# Classification: Bayes Decision Rule

## Error

### ◆ Error Probability



$$\begin{aligned}
 P(\text{error}) &= P(x \in \hat{X}_2, y_1) + P(x \in \hat{X}_1, y_2) \\
 &= P(x \in \hat{X}_2 | y_1)P(y_1) + P(x \in \hat{X}_1 | y_2)P(y_2) \\
 &= \int_{\hat{X}_2} p(x|y_1)P(y_1)dx + \int_{\hat{X}_1} p(x|y_2)P(y_2)dx
 \end{aligned}$$

$\hat{X}_1 = \{x | x \text{ is classified as } y_1\}$

$\hat{X}_2 = \{x | x \text{ is classified as } y_2\}$

**Prediction**

$X_1 = \{x | x \text{ belongs to } y_1\}$

$X_2 = \{x | x \text{ belongs to } y_2\}$

**True**

24

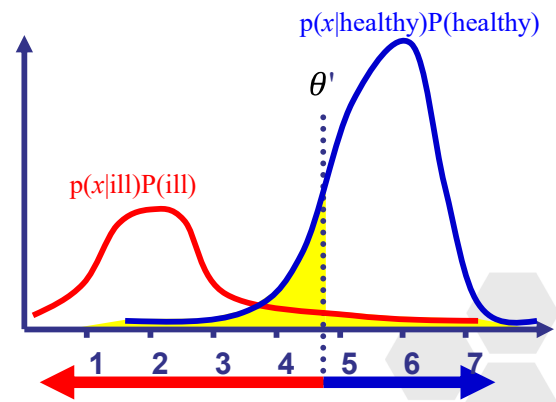
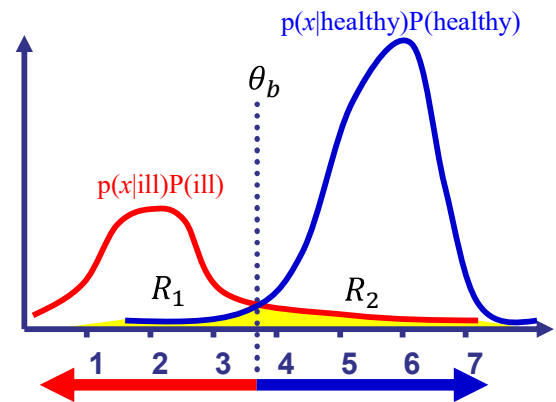
Dr. Patrick Chan @ SCUT

Artificial Intelligence III: Artificial Intelligence and Deep Learning - Lecture 3



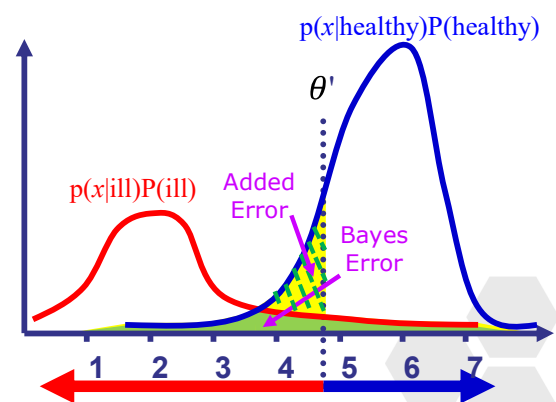
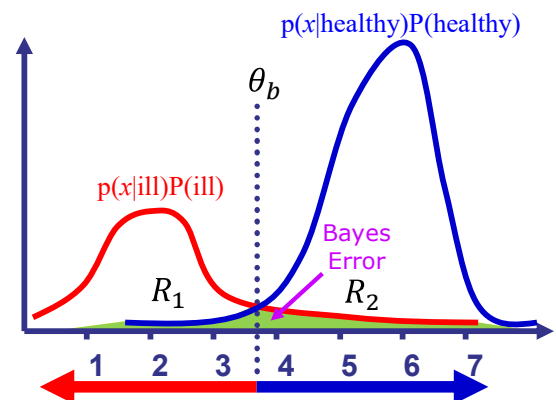
# Error

- ◆  $\theta_b$  or  $\theta'$  is better?
  - Error of  $\theta_b < \theta'$
  - $\theta_b$  is better
- ◆ Is any boundary better than Bayes Decision Rule ( $\theta_b$ )?
  - Bayes Rule is optimal (minimal classification error)



# Error

- ◆ **Error = Bayes Error + Added Error**
- ◆ **Bayes Error**  
Error of Bayes Rule
  - Cannot be reduced
  - Depend on the input space and application
- ◆ **Added Error**  
Additional error made by other classifiers
  - Can be reduced by selecting better parameters





# Extension to Multi-Class

- ◆ Extend to multi-class problem ( $c$  classes)

$$y = (y_1, y_2, \dots, y_c)$$

- ◆ Bayes Decision Rule

- $x$  is  $y_i$  if  $P(y_i | x)$  is maximum for  $i = 1 \dots c$

- ◆ Error for Bayes Decision Rule

$$P(\text{error} | x) = 1 - \max[ P(y_1 | x), P(y_2 | x), \dots, P(y_c | x) ]$$

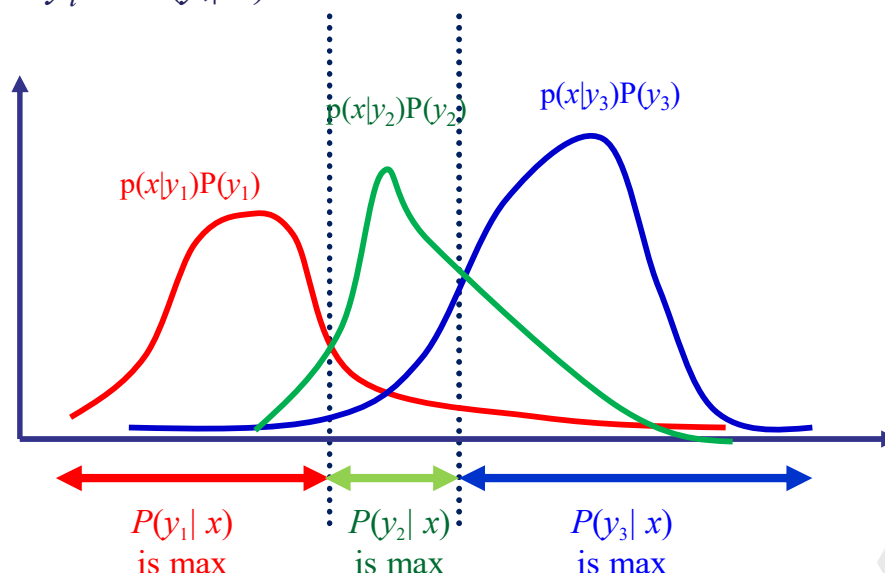


# Extension to Multi-Class

- ◆ Three-class example

- Bayes Decision Rule

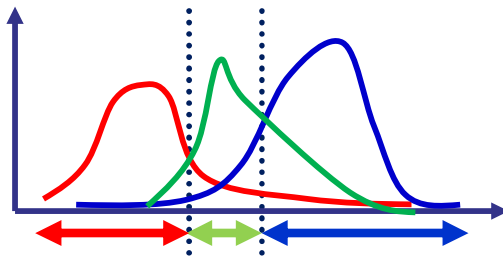
- $x$  is  $y_i$  if  $P(y_i | x)$  is max for  $i = 1 \dots 3$





# Classification: Bayes Decision Rule

## Extension to Multi-Class

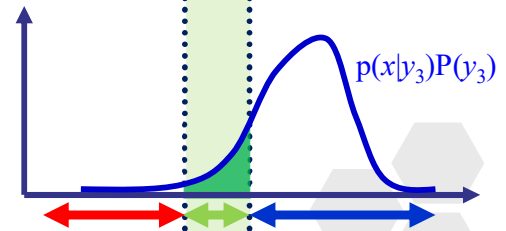
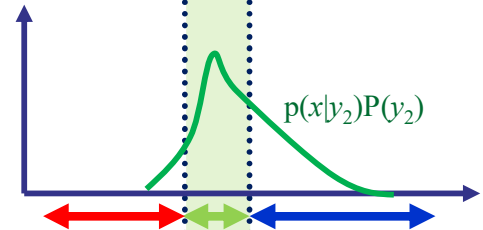
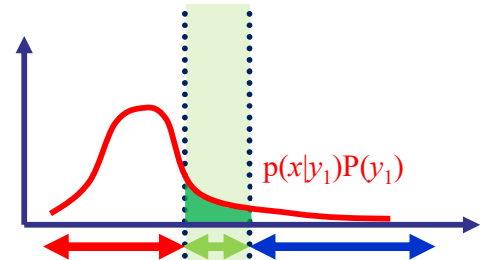


Error of Bayes Decision Rule:

$$P(\text{error} | x) = 1 - \max[P(y_1 | x), P(y_2 | x), P(y_3 | x)]$$

For example, in the **green region**,  
( $x$  is classified as  $y_2$  based on Bayes Rule)

$$\begin{aligned} P(\text{error} | x) &= P(y_1 | x) + P(y_3 | x) \quad * \text{Posterior probability (figures are not)} \\ &= 1 - P(y_2 | x) \\ &= 1 - \max_{i=1,2,3} P(y_i | x) \quad * \text{Must not be } P(y_2 | x) \end{aligned}$$



29

Dr. Patrick Chan @ SCUT

Artificial Intelligence III: Artificial Intelligence and Deep Learning - Lecture 3

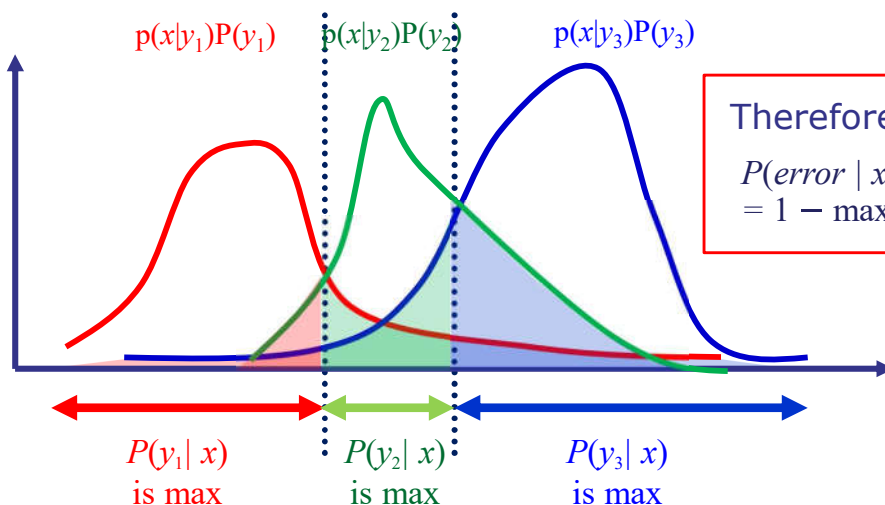


# Classification: Bayes Decision Rule

## Extension to Multi-Class

### ◆ Three-class example

#### ■ Error of Bayes Decision Rule



Therefore,

$$P(\text{error} | x) = 1 - \max[P(y_1 | x), P(y_2 | x), P(y_3 | x)]$$

$$\left. \begin{aligned} &P(y_1 | x) + P(y_3 | x) \\ &P(y_2 | x) + P(y_3 | x) \\ &P(y_1 | x) + P(y_2 | x) \end{aligned} \right\} P(\text{error} | x)$$

30

Dr. Patrick Chan @ SCUT

Artificial Intelligence III: Artificial Intelligence and Deep Learning - Lecture 3



# How to apply Bayes Rules?

## ◆ Bayes Rules:

If  $p(x|y_1)P(y_1) > p(x|y_2)P(y_2)$ ,  $x$  is classified as  $y_1$   
Otherwise, it is  $y_2$

## ◆ How to learn $p(x | y_i)$ and $P(y_i)$ in an application?

- $P(y_i)$  : Ratio of the class  $i$ 
  - $y_i$  is discrete, easy to estimate
- $p(x | y_i)$  : Distribution of samples in the class  $i$ 
  - $x$  is usually continues and has many dimensions, different to estimate



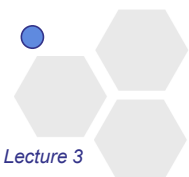
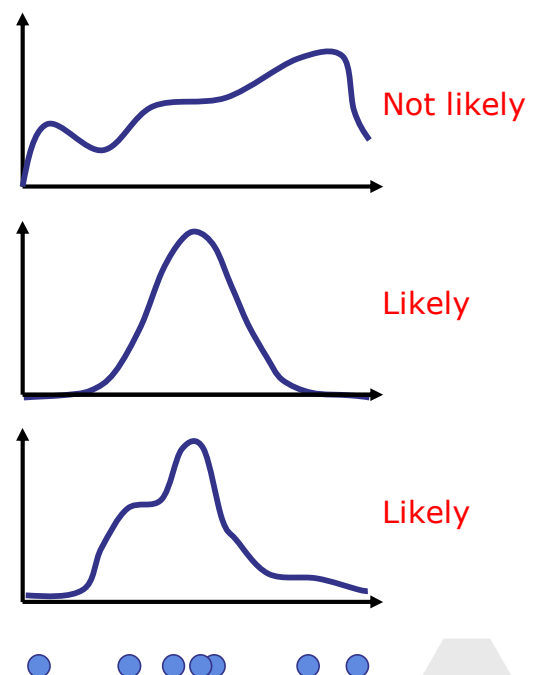
## How to apply Bayes Rules?

# $p(x | y_i)$ and $P(y_i)$ Estimation

## ◆ $p(x | y_i)$ means $p(x)$ and $x$ is from $y_i$

## ◆ Given samples of a class ( $x$ is from $y_i$ ), how can we know its real distribution $p(x)$ ?

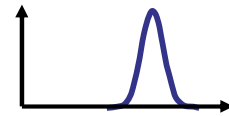
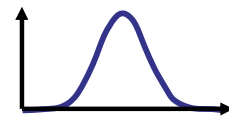
- By estimation



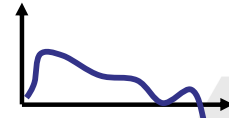
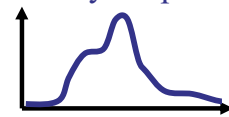


- ◆ How to estimate  $P(x | y_i)$ ?
  - **Parametric Methods** (Briefly Introduce here)
    - **Model-based Method** *Assume form of sample distribution (pdf) is known*
      - Estimate distribution parameters
      - Bias (Great if the assumptions are correct)
  - **Non-Parametric Methods** (Part 2, NN)
    - **Model Free Method** *No assumption on pdf*
      - A proper form for discriminant function is assumed
      - Usually sub-optimal, but good results generally

Shape is fixed



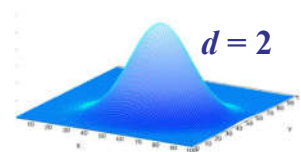
Any shape



- ◆ Assume **samples in each class follow normal distribution (Gaussian distribution),**

$$D \sim N(\mu, \sigma^2)$$

- ◆ **1-dimension:**  $p(x) = \frac{1}{\sqrt{2\pi}\sigma} \exp \left[ -\frac{1}{2} \left( \frac{x - \mu}{\sigma} \right)^2 \right]$



- ◆ **d-dimension:**  $p(x) = \frac{1}{(2\pi)^{d/2} |\Sigma|^{1/2}} \exp \left[ -\frac{1}{2} (x - \mu)^t \Sigma^{-1} (x - \mu) \right]$

$x = (x_1, x_2, \dots, x_d)^t$ : sample vector

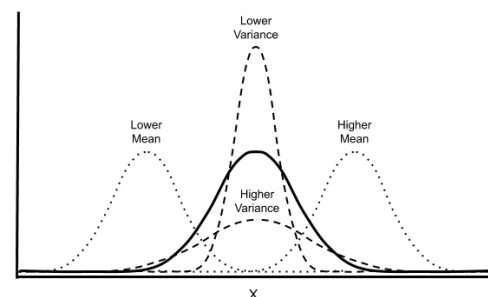
$\mu = (\mu_1, \mu_2, \dots, \mu_d)^t$ : mean vector

$\Sigma = d \times d$ : covariance matrix

$\sigma$ : variance (d=1 of covariance matrix)

$|\Sigma|$  and  $\Sigma^{-1}$ : determinant and inverse

$t$ : transpose





# Discriminant Functions for the Normal Density

- ◆ To simplify calculations by transforming multiplication into addition

$$p(x|y)P(y) \text{ and } p(x|y) = \frac{1}{(2\pi)^{d/2}|\Sigma|^{1/2}} \exp \left[ -\frac{1}{2} (x - \mu)^t \Sigma^{-1} (x - \mu) \right]$$

- ◆ The natural log function is a monotonic increasing function

$$\begin{aligned} p(x|y)P(y) &\propto \ln(p(x|y)P(y)) \\ &= \ln(p(x|y)) + \ln(P(y)) \end{aligned}$$

- ◆  $g(x)$  is used for comparison

$$g(x) = \ln(p(x|y)) + \ln(P(y))$$

35

Dr. Patrick Chan @ SCUT

Artificial Intelligence III: Artificial Intelligence and Deep Learning - Lecture 3



# Discriminant Functions for the Normal Density

- ◆ Substitute  $p(x|y)$  to  $g(x)$

$$\begin{aligned} g(x) &= \ln(p(x|y)) + \ln(P(y)) \\ p(x|y) &= \frac{1}{(2\pi)^{d/2}|\Sigma|^{1/2}} \exp \left[ -\frac{1}{2} (x - \mu)^t \Sigma^{-1} (x - \mu) \right] \end{aligned}$$

- ◆ Therefore:

$$g(x) = -\frac{d}{2} \ln 2\pi - \frac{1}{2} \ln |\Sigma| - \frac{1}{2} (x - \mu)^t \Sigma^{-1} (x - \mu) + \ln P(y)$$

36

Dr. Patrick Chan @ SCUT

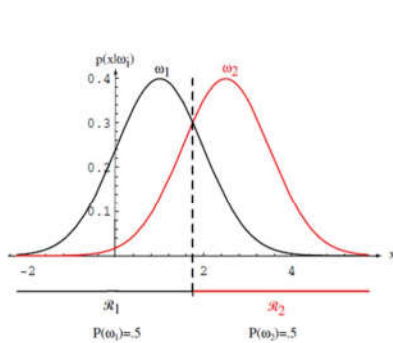
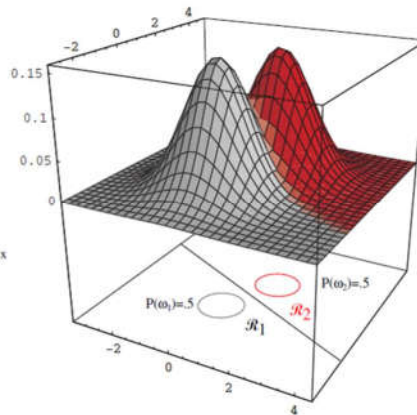
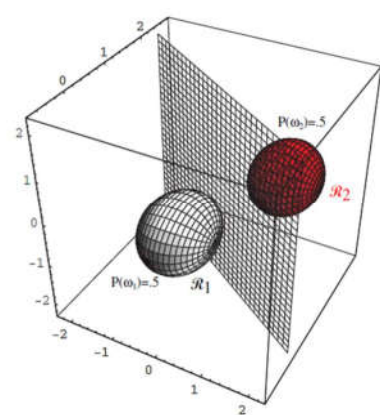
Artificial Intelligence III: Artificial Intelligence and Deep Learning - Lecture 3





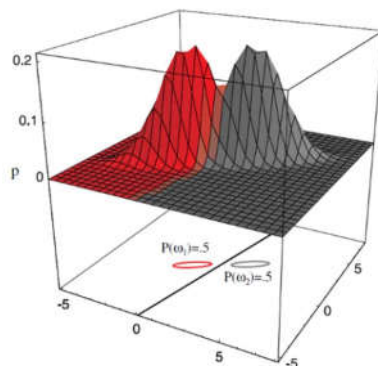
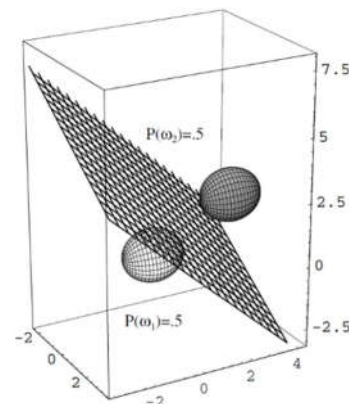
# Normal Distribution ( $\Sigma_i = \sigma^2 I$ )

- Assume the covariance matrices are the identity matrix, the distributions are spherical

 $d = 1$  $d = 2$  $d = 3$ 

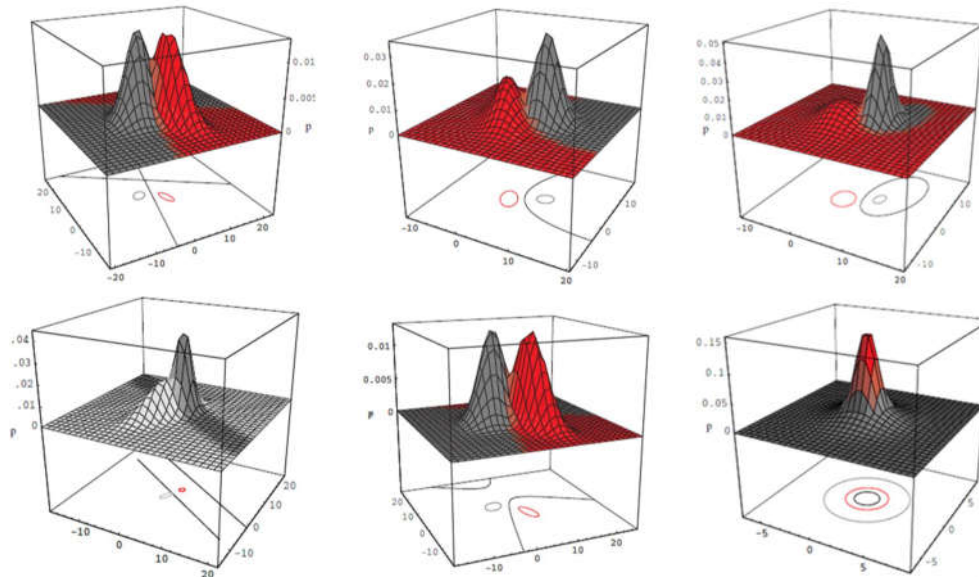
# Normal Distribution ( $\Sigma_i = \sigma^2 I$ )

- Assume each class has the same covariance matrix, the distributions is in ellipse sharp

 $d = 2$  $d = 3$

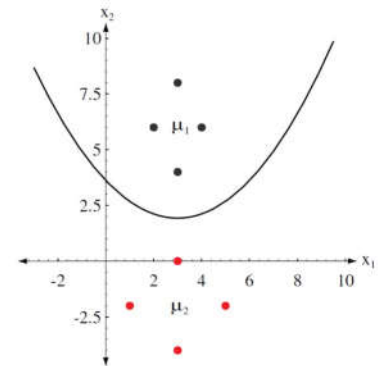


- ◆ The covariance matrix can be anything



## Discriminant Functions for Multivariate Normal Density

- ◆ Given  $\mu_1 = \begin{bmatrix} 3 \\ 6 \end{bmatrix}$   $\mu_2 = \begin{bmatrix} 3 \\ -2 \end{bmatrix}$   
 $\Sigma_1 = \begin{pmatrix} 1/2 & 0 \\ 0 & 2 \end{pmatrix}$   $\Sigma_2 = \begin{pmatrix} 2 & 0 \\ 0 & 2 \end{pmatrix}$
- ◆ Inverse  $\Sigma_1^{-1} = \begin{pmatrix} 2 & 0 \\ 0 & 1/2 \end{pmatrix}$   $\Sigma_2^{-1} = \begin{pmatrix} 1/2 & 0 \\ 0 & 1/2 \end{pmatrix}$
- ◆ Assume  $P(y_1) = P(y_2) = 0.5$
- ◆ Decision Boundary  $g_1(x) = -(x_1 - 3)^2 - \frac{1}{4}(x_2 - 6)^2 - \ln(4\pi)$   
 $g_2(x) = -\frac{1}{4}(x_1 - 3)^2 - \frac{1}{4}(x_2 + 2)^2 - \ln(8\pi)$   
 $g_1(x) - g_2(x) = -\frac{3}{4}(x_1 - 3)^2 - \frac{1}{2}(x_2 - 2)^2 - \ln(2) = 0$



$$g_i(x) = -\frac{1}{2}(x - \mu_i)^t \Sigma_i^{-1}(x - \mu_i) - \frac{d}{2} \ln 2\pi - \frac{1}{2} \ln |\Sigma_i| + \ln P(\omega_i)$$



## ◆ Multi-class problem

- Even with small number of classes, the shapes of the boundary regions is complex

