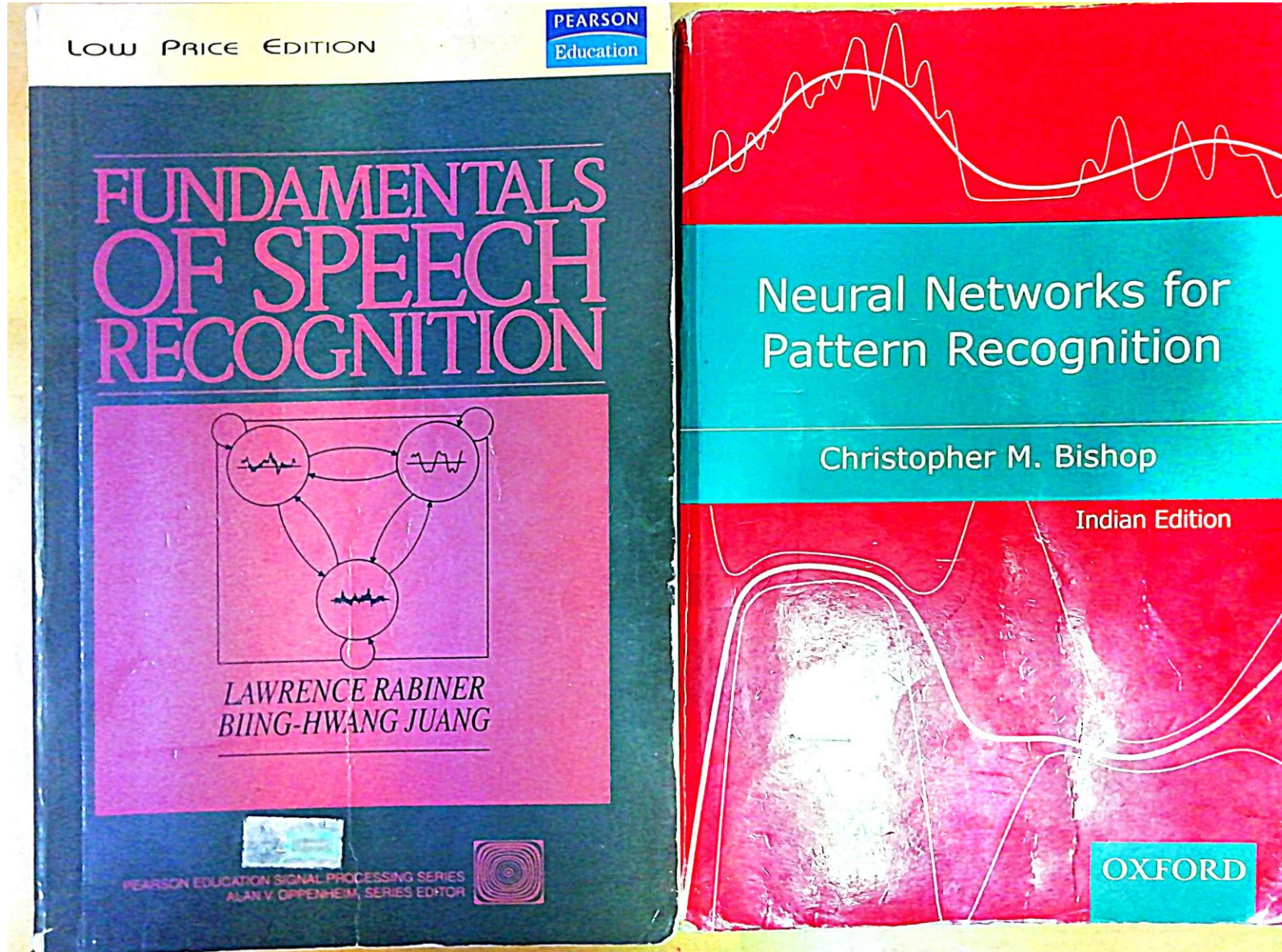


Statistical Pattern Recognition

The slide features a white background with the title 'Statistical Pattern Recognition' in a blue sans-serif font. On the right side, there is a large, solid blue shape that resembles a stylized wave or a bell curve. At the bottom left, there is a smaller blue shape with a series of white curved lines.

Books



Overview

- ❑ Pattern Classification
- ❑ Preprocessing & Feature Extraction
- ❑ Curse of Dimensionality
- ❑ Classification & Regression
- ❑ Estimation of Mapping Function
- ❑ Generalization
- ❑ Model Complexity (Bias & Variance)
- ❑ Regularization
- ❑ Baye's Theorem (Baye's Rule)
- ❑ Inference & Decision
- ❑ Decision Boundaries
- ❑ Probability Density Estimation



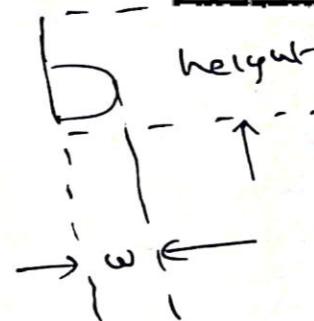
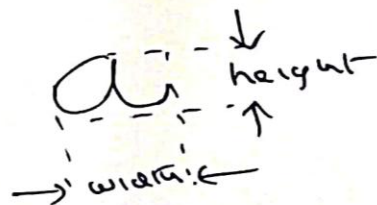
Pattern Classification

Pattern classification (Ex: character recognition)

$\boxed{a}$ $\boxed{b}$ $\Rightarrow$ Images 256×256

Total possible Images = $2^{8 \times 256 \times 256} \approx 10^{158000}$

Feature $\Rightarrow \frac{\text{height}}{\text{width}} = \tilde{x}_1$



Preprocessing

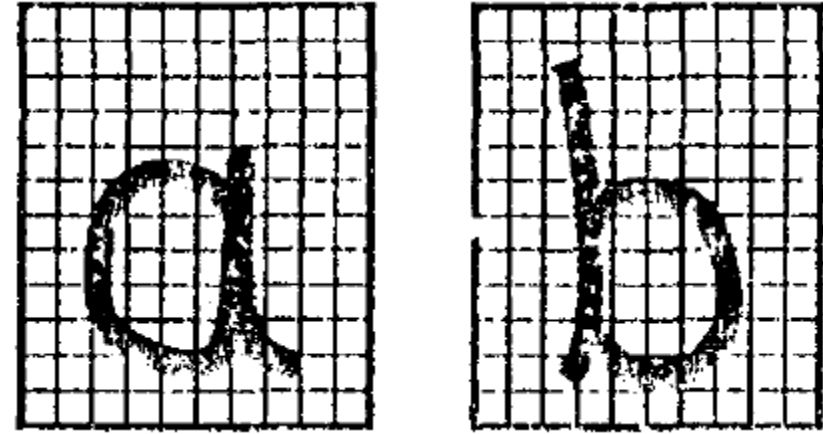
Raw data $\longrightarrow$ Features
(Pixels) $\left(\frac{\text{height}}{\text{width}}, \tilde{x}_1 \right)$

Overlapping of features $\Rightarrow$ Misclassification
Point of Intersection of Two Histograms $\Rightarrow$ Minimum Error in Classification

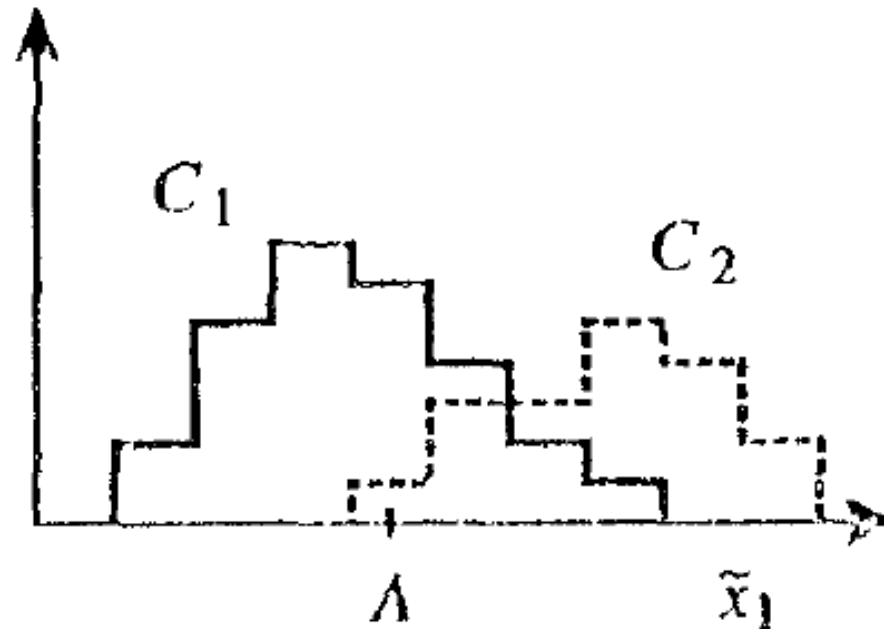


Pattern Classification (Cont..)

- ❑ Preprocessing
 - ✓ Raw data $\rightarrow$ Features
 - ✓ Pixel values $\rightarrow$ Height/width (x1)
 - ✓ Scale & Translation invariance

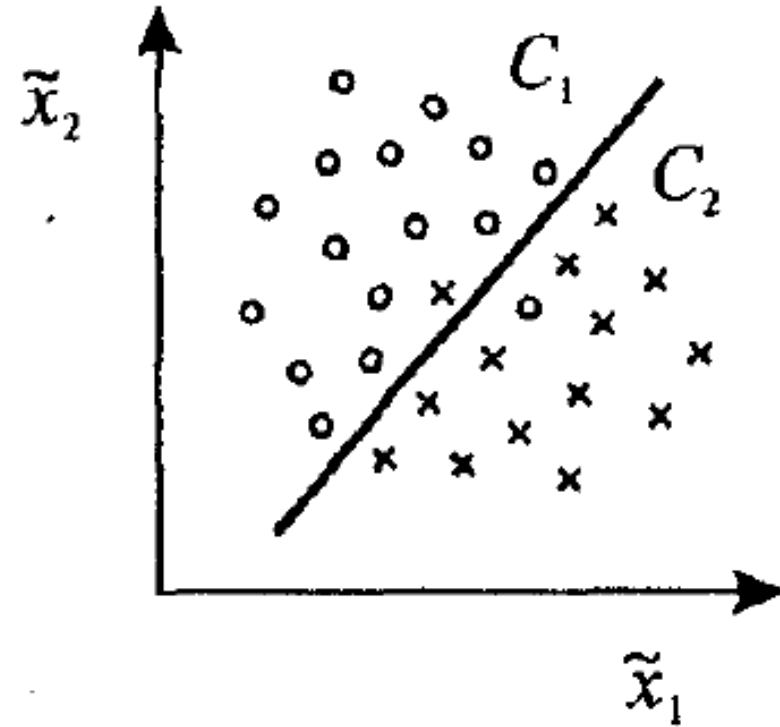


- ❑ Overlapping of features $\rightarrow$ Misclassification
- ❑ Point of intersection of two histograms $\rightarrow$ Minimum error in misclassification



Pattern Classification (Cont..)

- ❑ Preprocessing
 - ✓ Raw data $\rightarrow$ Features
 - ✓ Pixel values $\rightarrow$ Height/width (x1)
- ❑ Overlapping of features $\rightarrow$ Misclassification
- ❑ Point of intersection of two histograms $\rightarrow$ Minimum error in misclassification
- ❑ Increase in number of features $\rightarrow$ Reduce misclassification



Curse of Dimensionality

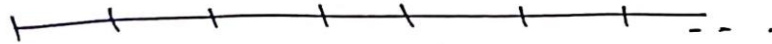
Curse of Dimensionality

d -features $\Rightarrow d$ -dimensional Feature Vector

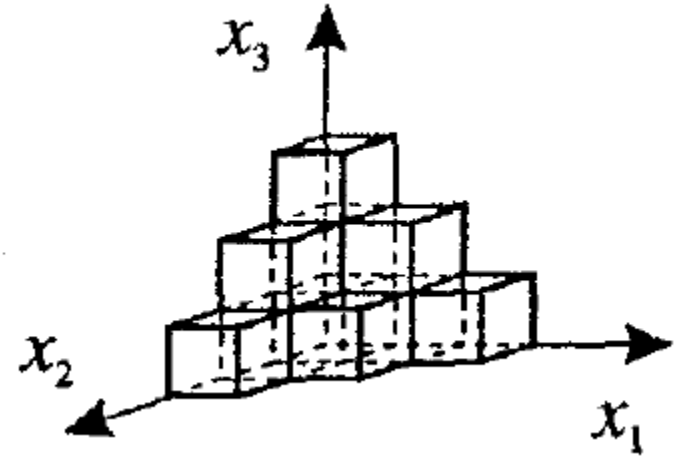
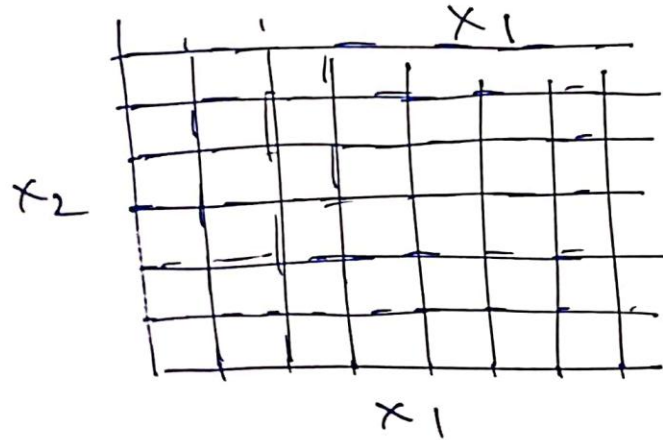
Each dimension (Feature) $\rightarrow M$ divisions

Cells $\Rightarrow M^d$

$d=1 \Rightarrow$
(M)

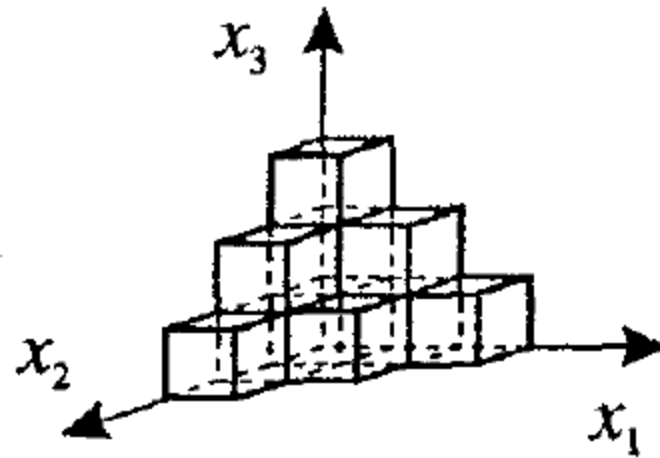


$d=2 \Rightarrow$
(M^2)



Curse of Dimensionality (Cont..)

- Concept of intrinsic Dimensionality
 - ✓ Input data points are correlated and restricted to subspace of lower dimensionality
- Interpolation
 - ✓ Output varies smoothly as a function of input variables
- Dimensionality Reduction
 - ✓ Limited & fixed size datasets



Classification & Regression

❑ Function Approximation

- ✓ Classification (Probabilities of memberships of different classes)
- ✓ Regression (Function defined in terms of average over random quantity)

❑ Classification : Discrete Classes

- ✓ Speech Recognition : 40 Classes (Phoneme as a class)
- ✓ Character Recognition : 26 Classes (Each alphabet as a class)

❑ Regression : Continuous Output

- ✓ Stock price prediction
- ✓ Weather prediction

Estimation of Mapping Function

❑ Function : The underlying function is modeled using polynomial curve fitting

$$y(x) = w_0 + w_1 x + w_2 x^2 + \dots + w_m x^m$$

$$y = f(x; w)$$

❑ Input : $x_1, x_2, \dots, x_n$

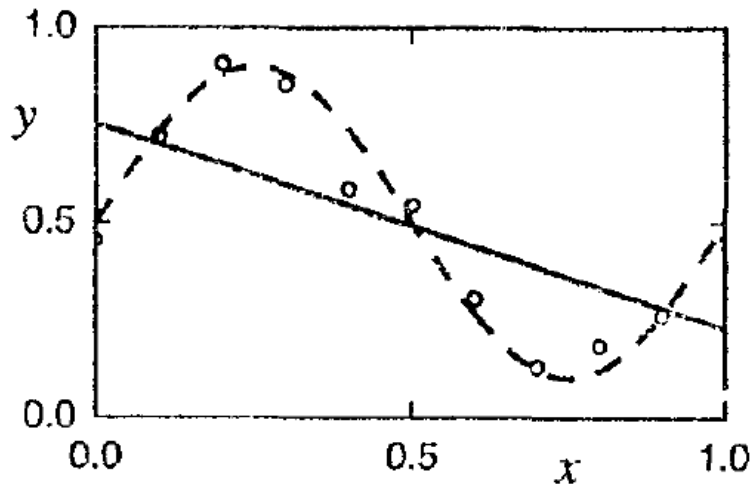
❑ Output (Target) : $t_1, t_2, \dots, t_n$

$$\text{❑ } h(x) = 0.5 + 0.4 \cdot \sin(2\pi x)$$

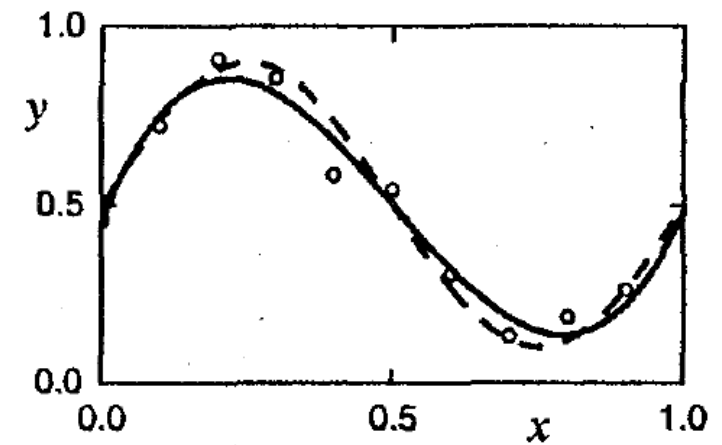
❑ $X_n \rightarrow h(x) + \text{Random noise}$

$$\text{❑ } E^{rms} = \sqrt{\sum_{n=1}^N (y(x_n; w^*) - t_n)^2}$$

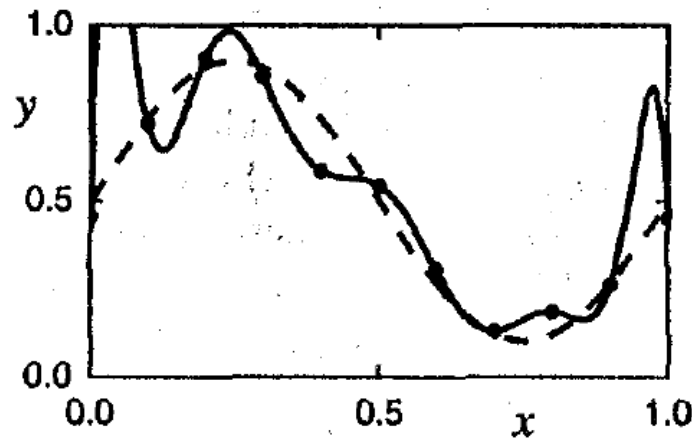
Estimation of Mapping Function



Simple Model ($M=1$)
Linear function



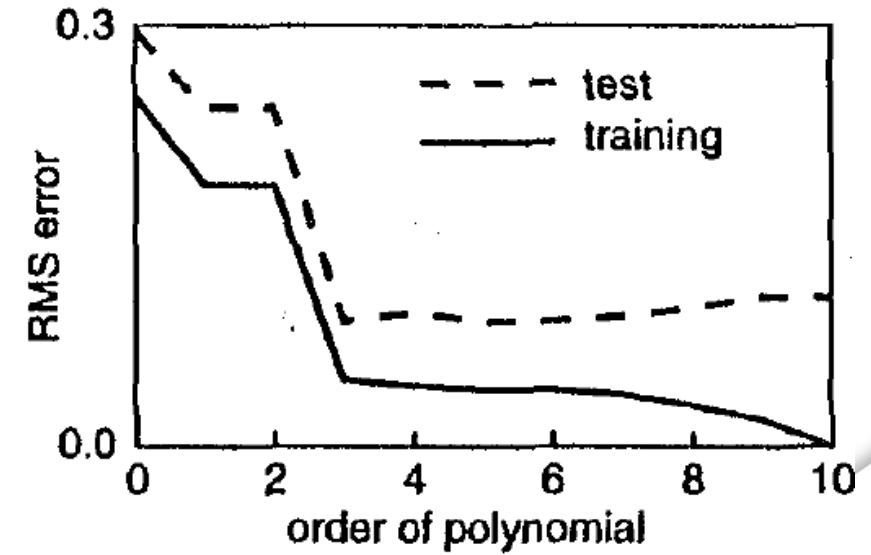
Optimal Model ($M=3$)
Cubic function



Complex Model ($M=10$)
10th Order Non-linear function

Estimation of Mapping Function (Cont..)

- ❑ Generalization
- ❑ Degrees of freedom $\rightarrow$ # free variables ($w_0, w_1, w_2, \dots$)
- ❑ Lower degrees of freedom $\rightarrow$ Higher Bias (under fitting)
- ❑ Higher degrees of freedom $\rightarrow$ Higher variance (over fitting, noisy)
- ❑ Generalization : Trade-off between Bias & Variance
- ❑ Good generalization : Low Bias & Low Variance



Model Complexity : Generalization & Regularization

□ Best Generalization

- ✓ Complexity of model is neither too small not too large
- ✓ Optimal complexity

□ Optimum Generalization → Controlling the effective complexity

- ✓ Polynomial fit with $M = 1$ (simple model with poor fit to data)
- ✓ $M = 3 \rightarrow$ Cubic polynomial (Optimal model with best generalization)
- ✓ $M = 10 \rightarrow$ Complex model over-fits the data

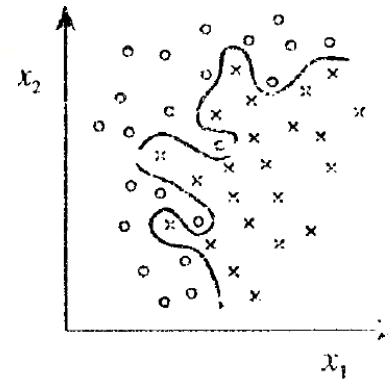
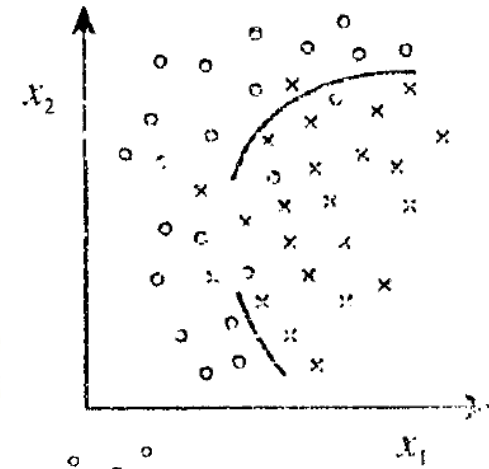
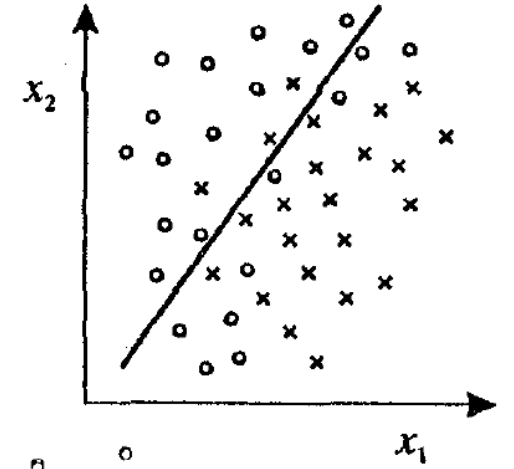
□ Effective complexity → Adding penalty term to error function

Effective complexity $\Rightarrow$ Adding Penalty to Error function

$$\tilde{E} = E + \lambda R; \quad R \rightarrow \text{Regularization Term} \\ \rightarrow \text{depends on mapping fun.}$$

$$R = \frac{1}{2} \int \left(\frac{d^2 y}{dx^2} \right)^2 dx$$

Large variance (noise) $\Rightarrow R \uparrow$



Baye's Theorem (Baye's Rule)

Baye's Theorem

$C_k \rightarrow k^{\text{th}} \text{ class}$

$x \rightarrow \text{Data vector}$

c_1	.	.	.	.	.	.	.		(70)
c_2				.	.	.	.	.	(30)
	x_e								

$$P(C_k, x) \Rightarrow \text{Joint probability of } x \text{ \& } C_k$$

$$= P(C_k | x) P(x) = P(x | C_k) P(C_k)$$

$$P(c_1) = \frac{70}{100} ; \quad P(c_2) = \frac{30}{100}$$

$$P(C_k, x_e) < \begin{cases} P(c_1, x_e) = \frac{7}{100} \\ P(c_2, x_e) = \frac{3}{100} \end{cases}$$

$$P(c_1, x_e) = P(x_e | c_1) P(c_1) = \frac{7}{70} \cdot \frac{70}{100} = \frac{7}{100}$$

$$P(\cancel{c_2}, x_e) = P(c_1 | x_e) P(x_e) = \frac{7}{10} \cdot \frac{10}{100} = \frac{7}{100}$$

$$P(c_2, x_e) = P(c_2 | x_e) \cdot P(x_e) = \frac{3}{10} \cdot \frac{10}{100} = \frac{3}{100}$$

$$= P(x_e | c_2) P(c_2) = \frac{3}{30} \cdot \frac{30}{100} = \frac{3}{100}$$

Baye's Theorem (Baye's Rule) Cont..

$$P(c_k, x_e) = P(c_k | x_e) P(x_e) = P(x_e | c_k) P(c_k)$$

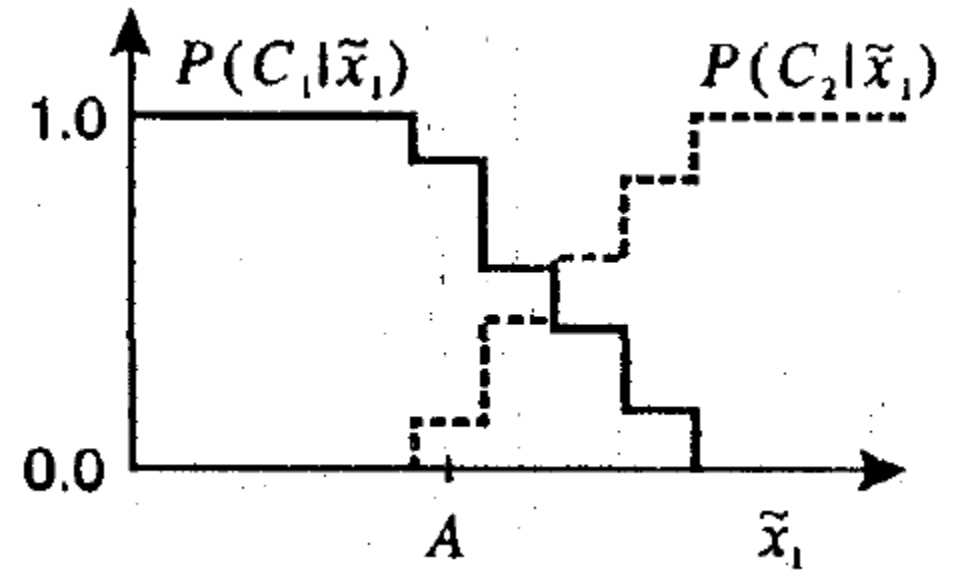
$$P(c_k | x_e) = \frac{P(x_e | c_k) \cdot P(c_k)}{P(x_e)}$$

$P(c_k | x_e) \rightarrow$ Posterior probability

$P(x_e | c_k) \rightarrow$ class conditional probability

$P(c_k) \rightarrow$ Prior probability

$P(x_e) \rightarrow$ Normalization factor
(unconditional probability)



Baye's Theorem (Baye's Rule) Cont..

$P(c_k | x_e)$ — posterior probability

$$\sum_{k=1}^K P(c_k | x_e) = 1$$

$$\sum_{k=1}^K \frac{P(x_e | c_k) P(c_k)}{P(x_e)} = 1$$

$$\Rightarrow \sum_{k=1}^K P(x_e | c_k) P(c_k) = P(x_e)$$

Two-class problem

$$P(c_1 | x_e) + P(c_2 | x_e) = 1$$

$$P(x_e) = P(x_e | c_1) P(c_1) + P(x_e | c_2) P(c_2)$$

Inference & Decision

Inference $\rightarrow$ Estimation of posterior probabilities

$$P(c_k | x) = \frac{P(x | c_k) P(c_k)}{P(x)}$$

$P(x | c_k) \Rightarrow$ Prob density estimation of a specific class

$P(c_k) \Rightarrow$ Prior probability estimated from general statistics

$P(x) \Rightarrow$ Normalization factor

Decision Rule $\Rightarrow$ Minimum Probability of Misclassification

Bayes' Rule $\Rightarrow$ class label assign to Max posterior prob

$$\text{Label} = c_k \Rightarrow P(c_k | x) > P(c_j | x) \quad \forall j \neq k$$

$$\text{Posterior Prob} = \frac{\text{Likelihood} \times \text{Prior Prob}}{\text{Normalized factor}}$$

Decision Boundaries

Decision towards C_k if $P(C_k|x) > P(C_j|x) \quad \forall j \neq k$

$$\begin{aligned} P_e &= P(x \in R_2, c_1) + P(x \in R_1, c_2) \\ &= P(x \in R_2 | c_1) P(c_1) + P(x \in R_1 | c_2) P(c_2) \\ &= \int_{R_2} P(x|c_1) P(c_1) dx + \int_{R_1} P(x|c_2) P(c_2) dx \end{aligned}$$

$$\begin{aligned} P_{\text{correct}} &= \sum_k P(x \in R_k, c_k) \\ &= \sum_k P(x \in R_k | c_k) P(c_k) \\ &= \sum_k \int_{R_k} P(x|c_k) P(c_k) dx \end{aligned}$$

