

Probability Density Estimation



Overview

❑ Parametric Methods

- ✓ Maximum Likelihood Estimation Method
- ✓ Bayesian Inference Method

❑ Non-parametric Methods

- ✓ Histogram
- ✓ Kernel-based Approaches
- ✓ K-Nearest Neighbor (KNN) Approach

❑ Kullback-Leibler (KL) Distance

❑ Semi-Parametric Methods : Mixture Models (Gaussian Mixture Models (GMMs))

- ✓ Maximum Likelihood (Nonlinear Optimization)
- ✓ The EM Algorithm



Overview (Cont..)

❑ Parametric Methods

- ✓ Specific functional form is chosen
- ✓ Evaluation with new data is faster

❑ Non-parametric Methods

- ✓ No specific functional form
- ✓ Number of parameters grows with number of data points
- ✓ Very slow for evaluation of new data

❑ Semi-Parametric Methods

- ✓ No specific functional form
- ✓ Number of parameters grows with complexity
- ✓ Computationally intensive



Parametric Methods

Ex: Normal & Gaussian distribution

Estimator of parameter of a Model

- Maximum likelihood estimation Method
- Bayesian Inference Method

Normal Distribution (Single Variable)

$$P(x) = \frac{1}{(2\pi\sigma^2)^{\frac{1}{2}}} \exp\left[-\frac{(x-\mu)^2}{2\sigma^2}\right]; \int_{-\infty}^{\infty} p(x) = 1$$

μ - Mean; σ = Standard Deviation; σ^2 = Variance

$$\mu = E[x] = \int_{-\infty}^{\infty} x p(x) dx \Rightarrow \text{Expectation of } x$$

$$\sigma^2 = E[(x-\mu)^2] = \int_{-\infty}^{\infty} (x-\mu)^2 p(x) dx$$

$E[\cdot]$ = Expectation

Parametric Methods (Cont..)

MULTIVARIATE NORMAL DISTRIBUTION

$$P(x) = \frac{1}{(2\pi)^{\frac{d}{2}} |\Sigma|^{\frac{1}{2}}} \exp \left[-\frac{1}{2} (x-\mu)^T \Sigma^{-1} (x-\mu) \right]$$

$\Sigma \rightarrow$ co-variance Matrix of size $d \times d$

$\mu \rightarrow$ Mean vector of size d

$$\mu = E(x) ; \Sigma = E[(x-\mu)(x-\mu)^T]$$

$\Sigma \rightarrow$ Symmetric Matrix

Independent parameter = (μ, Σ)

$$= d + \frac{d(d+1)}{2} = \frac{d(d+3)}{2}$$

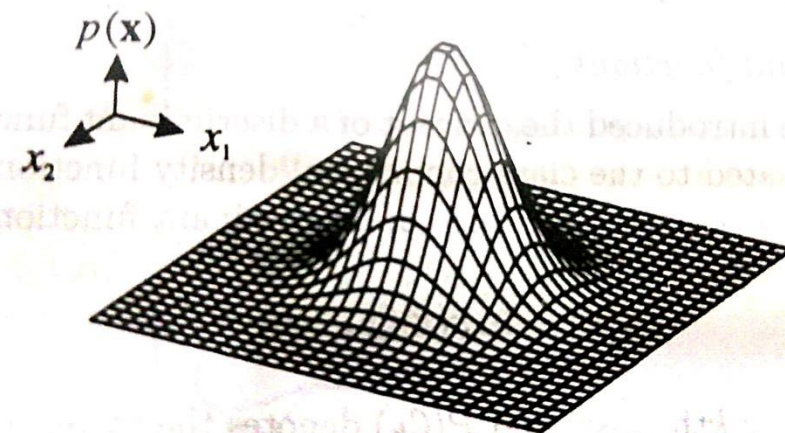
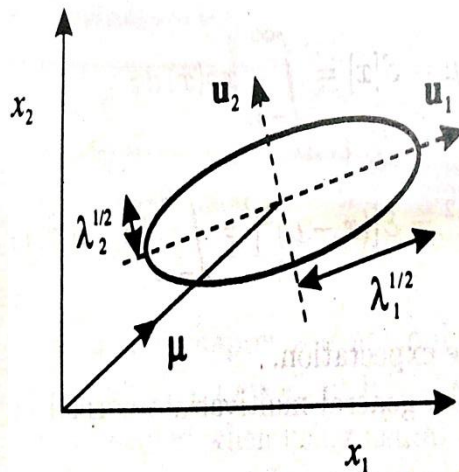
$$\Delta^2 = (x-\mu)^T \Sigma^{-1} (x-\mu) \rightarrow \text{Mahalanobis distance from } x \text{ to } \mu$$

Hyper-ellipsoids ($H E_1$)

Principal axes of $H E_1$ are Eigen vectors

$$\Sigma u_{\hat{n}} = \lambda_{\hat{n}} u_{\hat{n}}$$

$\lambda_{\hat{n}} \rightarrow$ Eigen values.



Simplification of parametric models

Components of x are statistically independent

$\Rightarrow \Sigma \rightarrow$ diagonal matrix

parameter = $(\mu, \Sigma) \Rightarrow 2d$

$$p(x) = \prod_{\hat{n}=1}^d p(x_{\hat{n}})$$

$$\sigma_{\hat{n}} = \sigma \quad \text{for all } \hat{n}$$

parameter = $d+1$

$\Rightarrow$ Hyper spheres

Parametric Methods (Cont..)

MAXIMUM LIKELIHOOD ESTIMATION

Optimum values of the parameter by maximizing the likelihood function derived from training data.

parameter set $\theta = [\theta_1, \theta_2, \dots, \theta_M]$

Data set $X = [x^1, x^2, \dots, x^N]$

$$p(x) = p(x|\theta)$$

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

$L(\theta)$ = likelihood of θ given X

$$\text{ML function: } p(X|\theta) = \prod_{n=1}^N p(x^n|\theta) = L(\theta)$$

$$\text{Error function } E = -\ln L(\theta) = \sum_{n=1}^N \ln p(x^n|\theta)$$

$$\frac{\partial E}{\partial \theta} = 0 \Rightarrow \text{Optimal parameter}$$

$$\hat{\mu} = \frac{1}{N} \sum_{n=1}^N x^n; \quad \hat{\sigma}^2 = \frac{1}{N} \sum_{n=1}^N (x^n - \hat{\mu})^T (x^n - \hat{\mu})$$

Parametric Methods (Cont..)

BAYESIAN INFERENCE

Parameters are represented by pdfs

Prior pdfs are assumed

Posterior parameter pdfs are estimated by observing the dataset.

$$p(x|x) = \int p(x, \theta|x) d\theta$$

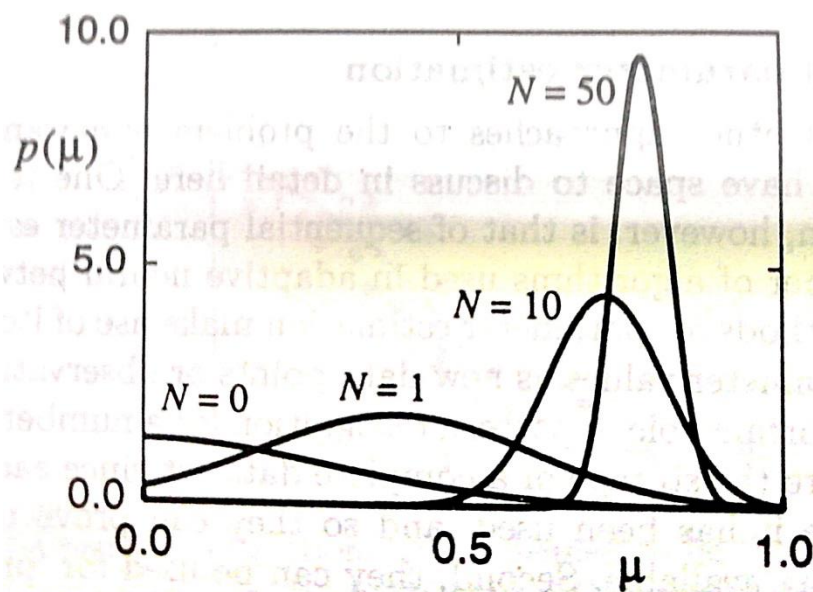
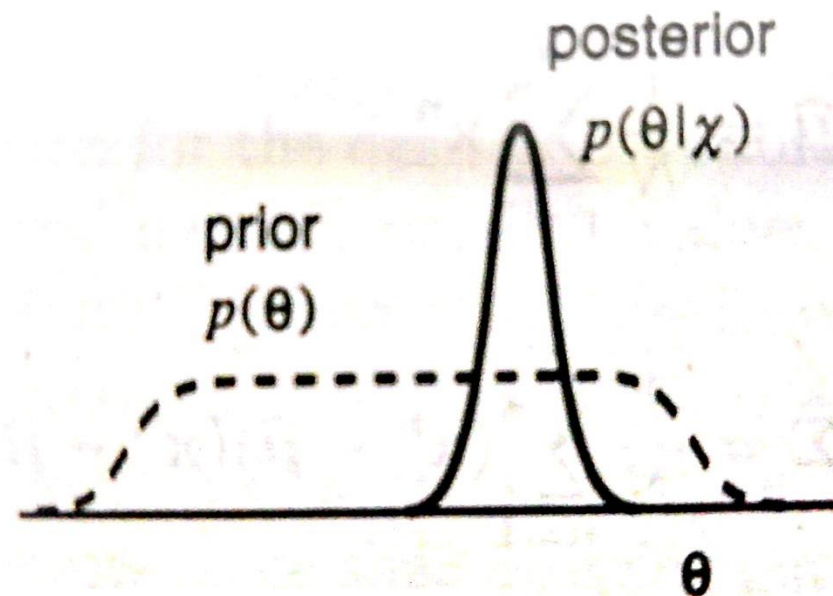
$$p(x, \theta|x) = p(x|\theta, x) p(\theta|x)$$

As $p(x|\theta, x)$ is independent of x

$$p(x, \theta|x) = p(x|\theta) p(\theta|x)$$

$$p(x|x) = \int p(x|\theta) p(\theta|x) d\theta$$

pdf of x is weighted average over all posterior probabilities of θ .



Parametric Methods (Cont..)

$$p(\theta|x) = \frac{p(x|\theta) p(\theta)}{p(x)} = \frac{\prod_{n=1}^N p(x^n|\theta) p(\theta)}{p(x)} = \frac{p(\theta)}{p(x)} \prod_{n=1}^N p(x^n|\theta)$$

Conjugate Prior: $\{p(x|\theta), p(\theta)\} \Rightarrow p(\theta|x)$

$$\text{initial pdf of } \mu = p_0(\mu) = \frac{1}{(2\pi\sigma_0^2)^{1/2}} \exp\left[-\frac{(\mu - \mu_0)^2}{2\sigma_0^2}\right]$$

$$\text{Posterior pdf of } \mu = p_N(\mu|x) = \frac{p_0(\mu)}{p(x)} \prod_{n=1}^N p(x^n|\mu)$$

$$\mu_N = \frac{N\sigma_0^2}{N\sigma_0^2 + \sigma^2} \bar{x} + \frac{\sigma^2}{N\sigma_0^2 + \sigma^2} \mu_0$$

$$\frac{1}{\sigma_N^2} = \frac{N}{\sigma^2} + \frac{1}{\sigma_0^2}$$

$$\bar{x} = \frac{1}{N} \sum_{n=1}^N x_n$$

Non-Parametric Methods

- ❑ Histograms

- ❑ Kernel-based approaches

- ❑ K-Nearest Neighbor (KNN)

- ❑ Histograms (Simple Models)

- Role of smoothing parameters

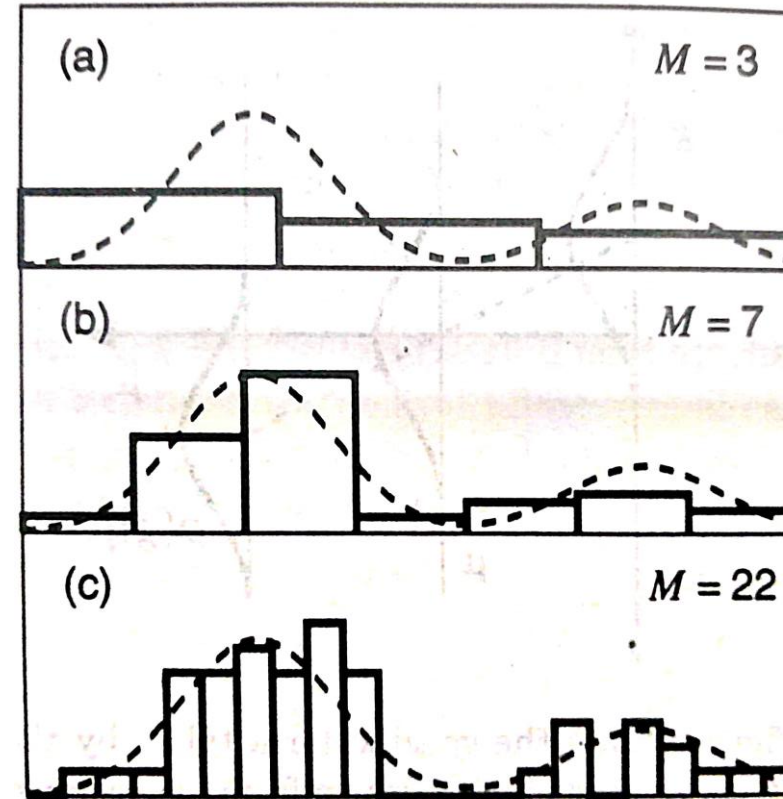
- ✓ Smoothing parameter $\rightarrow$ Number of bins
- ✓ Lower number of bins (smooth out crucial details)
- ✓ Higher number of bins (spiky & noisy)

- Advantages

- ✓ Can be constructed sequentially, entire data is not required at one shot. No storage issues.

- Disadvantages

- ✓ Estimated density is not smooth
- ✓ Curse of dimensionality with high-dimensionality data



Non-Parametric Methods : Density Estimation

Density Estimation in General

P - prob of vector x drawn from unknown density $p(x)$ will fall inside some region R of x -space is

$$P = \int_R p(x') dx'$$

Prob of K points fall within region R out of N points drawn

$$P_R(K) = \frac{N!}{K!(N-K)!} P^K (1-P)^{N-K} \Rightarrow \text{Binomial law}$$

Mean of K points fall in these region is given by

$$\text{Mean} = E \left[\frac{K}{N} \right] = P \approx \frac{K}{N}$$

$$\text{Variance} = E \left[\left(\frac{K}{N} - P \right)^2 \right] = \frac{P(1-P)}{N}$$

if $p(x)$ is continuous $\Rightarrow P = \int_R p(x') dx' \approx p(x) v$

$$\Rightarrow \frac{K}{N} = p(x) v \Rightarrow \boxed{p(x) = \frac{K}{Nv}}$$

Non-Parametric Methods : Density Estimation

Practical Density Estimation

$P \approx \frac{K}{N} \Rightarrow$ Region "R" to be relatively large $\Rightarrow P$ will be large.

$P \approx p(x) V \Rightarrow p(x) = \frac{P}{V} = \frac{K}{NV} \Rightarrow p(x)$ is approx constant when "R" is relatively small.

Choice of R for the best estimate of $p(x)$

K-Nearest Neighbor (KNN)

Fix the value of $K \Rightarrow$ Def the corresponding volume "V" from data

Kernel-based density

Fix the volume $V \Rightarrow$ Def "K" from data

For infinite "N" $\Rightarrow$ Both KNN & Kernel based methods will converge.

Non-Parametric Methods : Kernel-based Method

Kernel Based Method

Region "R" to be hypercube with sides of length "h" centered on the point x in d-dimensional space

Volume of the hypercube $V = h^d$

Kernel function $H(u) = \begin{cases} 1 & |u_j| < \frac{1}{2} \quad j = 1, 2, \dots, d. \\ 0 & \text{otherwise} \end{cases}$

$H\left[\frac{x - x^m}{h}\right] = 1$; if x^m falls within hypercube
 $= 0$ otherwise

$$K = \sum_{m=1}^N H\left[\frac{x - x^m}{h}\right]$$

$$\hat{p}(x) = \frac{K}{NV} = \frac{1}{N} \sum_{m=1}^N \frac{1}{h^d} H\left[\frac{x - x^m}{h}\right]$$

Estimated Density = Superposition of N cubes of side "h", with each cube centered on one of the data points
Cubes vs bins in case of Histogram

Non-Parametric Methods : Kernel-based Method

Kernel-Based Methods (cont...)

Smoothing the density by choosing different forms of kernel functions $H(u)$

Multivariate Normal (Gaussian) kernel

$$\tilde{p}(x) = \frac{1}{N} \sum_{n=1}^N \frac{1}{(2\pi h^2)^{d/2}} \exp \left[-\frac{\|x - x_n\|^2}{2h^2} \right]$$

$$E[\tilde{p}(x)] = \frac{1}{N} \sum_{n=1}^N E \left[\frac{1}{h^d} H \left(\frac{x - x_n}{h} \right) \right]$$

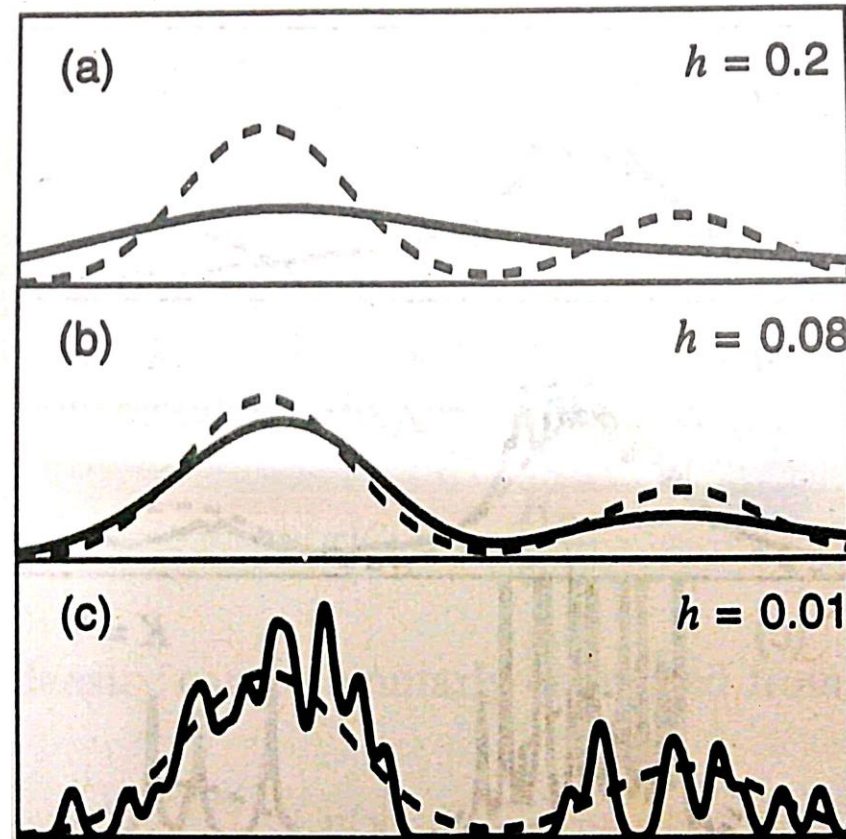
$$= E \left[\frac{1}{h^d} H \left(\frac{x - x'}{h} \right) \right]$$

$$= \int \frac{1}{h^d} H \left(\frac{x - x'}{h} \right) p(x') dx'$$

Convolution of true density with kernel function
 $\Rightarrow$ Smoothed version of true density.

$h \rightarrow 0 \Rightarrow$ kernel f_{or} is delta f_{un} $\Rightarrow \tilde{p}(x) = p(x)$

For finite sample size, small values of h leads to noisy representation of $\tilde{p}(x)$



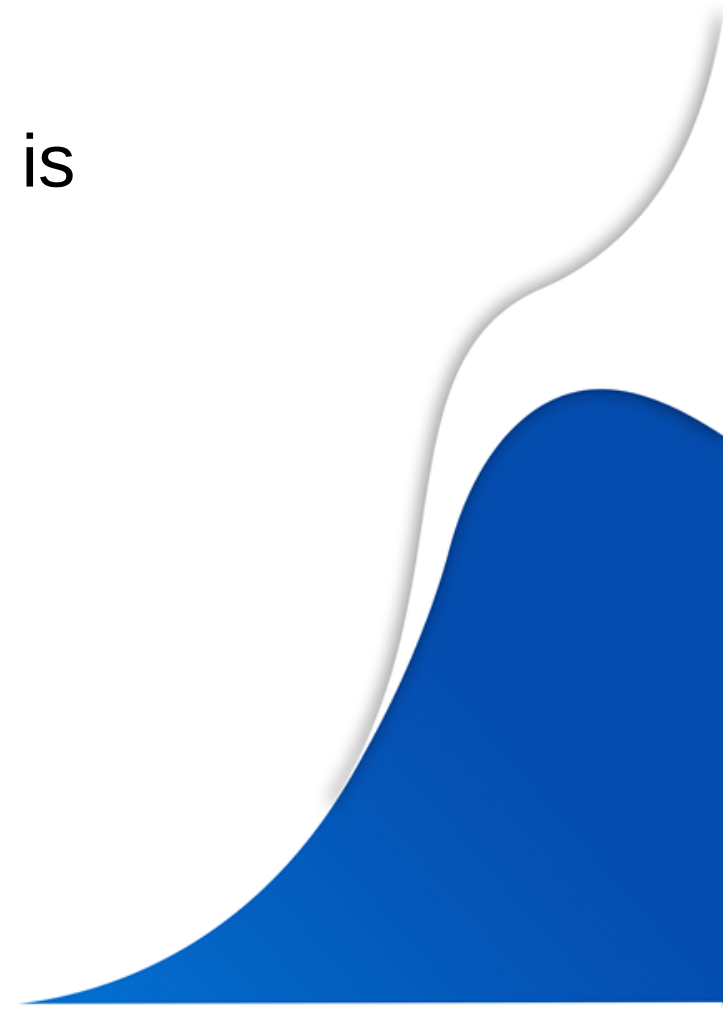
Non-Parametric Methods : Kernel-based Method

❑ Drawbacks

- ✓ All data points to be stored
- ✓ Evaluation of density is very slow if # data points is large
- ✓ It gives biased estimate of density

❑ Solution to fast density estimation

- ✓ Use of fewer kernel functions
- ✓ Adapt their positions and width as per the data



Non-Parametric Methods : K-Nearest Neighbor

- ❑ Problem with Kernel-based approach → Fixed “h” for all data points
- ❑ Remedy → KNN : Fix “K” value & allow “V” to vary
- ❑ Disadvantages of K-Nearest Neighbor (KNN)
 - ✓ Estimated PD is not true density, because integral over all X-space diverges.
 - ✓ All the training data points must be retained.
 - ✓ Large computer storage is required to store the whole data
 - ✓ Require large amount of processing to evaluate the density at new data points.

Non-Parametric Methods : KNN-Classifier

KNN: classifier

N - Total number of data points
 N_k - # points of class C_k

$$N = \sum_{k=1}^K N_k$$

For estimation of $p(x|C_k)$
 generate a hyper-sphere around " x " which contains " K " points
 Among " K " points, K_k corresponds to class C_k

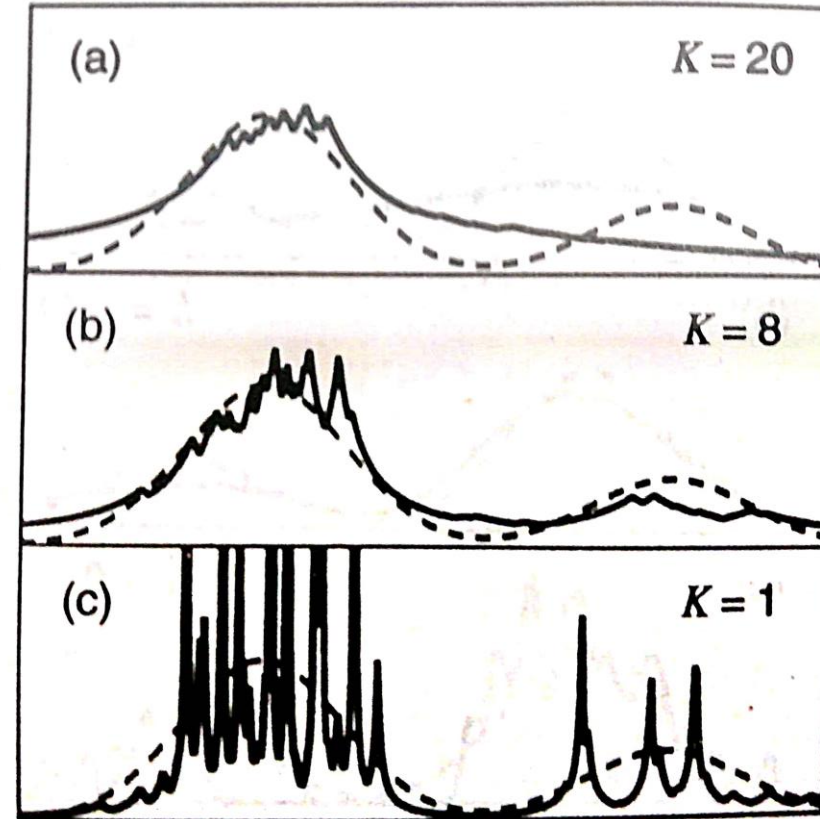
$$p(x|C_k) = \frac{K_k}{N_k V} \quad \text{— class conditional density}$$

$$p(x) = \frac{K}{NV} \quad \text{— un conditional density}$$

$$p(C_k) = \frac{N_k}{N}$$

$$p(C_k|x) = \frac{p(x|C_k) p(C_k)}{p(x)} = \frac{\frac{K_k}{N_k V} \cdot \frac{N_k}{N}}{\frac{K}{NV}} = \frac{K_k}{K}$$

To minimize the prob of misclassification, the vector x is assigned
 to class " C_k " for which the ratio of $\frac{K_k}{K}$ is largest
 $K=1 \Rightarrow$ Nearest Neighbor



Non-Parametric Methods : Smoothing Parameters

Smoothing Parameters

Histograms — width of the bins & # bins

Kernel Methods — width of the kernel "h"

kNN — value of K

over-smoothed — high bias — poor estimator

Insufficient smoothing — high variance — Noisy density —

Terms in a polynomial & # hidden units in NNs

$$L(h) = \frac{1}{N} \sum_{n=1}^N \log p(x^n | h; x^1, x^2, \dots, x^N)$$

$L(h)$ is maximized as $h \rightarrow 0$ (impulse at each data point)

$$E[-\log L(h)] = -\lim_{N \rightarrow \infty} \frac{1}{N} \sum_{n=1}^N \log \tilde{p}(x^n) = -\int p(x) \log \tilde{p}(x) dx$$

Distance bet true & estimated densities ~~is~~ $\int p(x) \log \frac{\tilde{p}(x)}{p(x)} dx$

$$L = \int p(x) \log p(x) - \int p(x) \log \tilde{p}(x) = -\int p(x) \log \frac{\tilde{p}(x)}{p(x)} dx$$

K-L distance : Kullback - Leibler distance

Distance is weighted by $p(x)$

Semi-Parametric Models : Mixture Models

❑ Parametric Methods

- ✓ Specific functional form is chosen
- ✓ Evaluation with new data is faster

❑ Non-parametric Methods

- ✓ No specific functional form
- ✓ Number of parameters grows with number of data points
- ✓ Very slow for evaluation of new data

❑ Semi-Parametric Methods

- ✓ No specific functional form
- ✓ Number of parameters grows with complexity
- ✓ Computationally intensive



Semi-Parametric Models : Mixture Models (Cont..)

Mixture Models (cont..)

$$p(x) = \sum_{j=1}^M p(x|j) p(j) \quad \text{— linear combination of component densities}$$

$$\sum_{j=1}^M p(j) = 1 \quad 0 \leq p(j) \leq 1$$

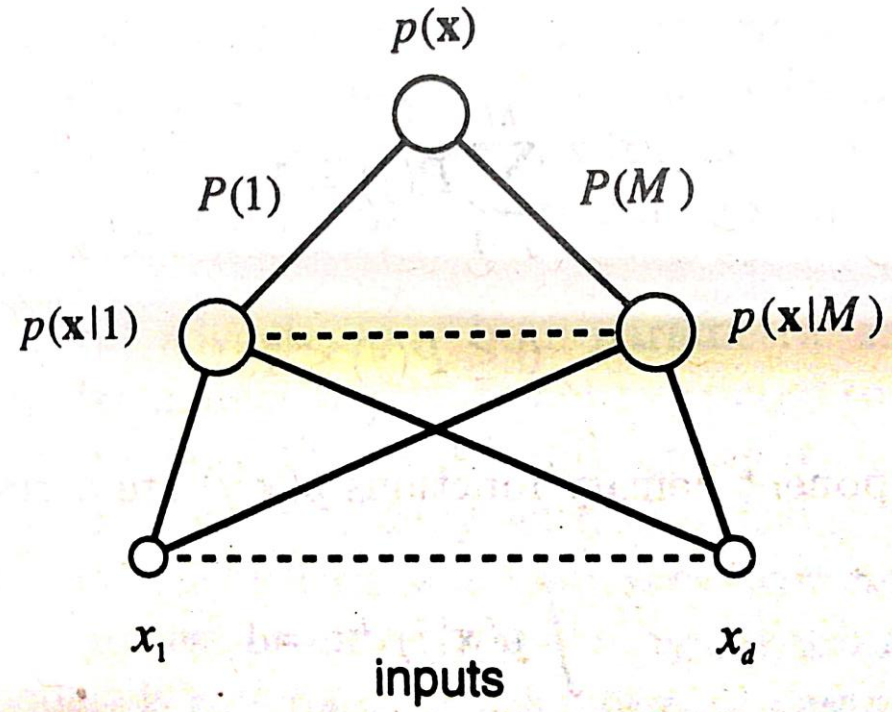
$$\int p(x|j) = 1$$

$p(x)$ — Mixture distribution

$p(j)$ — Mixing parameter (weight)

$p(x|j)$ — Component density

Property: For many choices of component density $f_{j=1}^M$, they can approximate any continuous density to arbitrary accuracy provided M is sufficiently large components & provided the parameters of the model are chosen



Semi-Parametric Models : Mixture Models (Cont..)

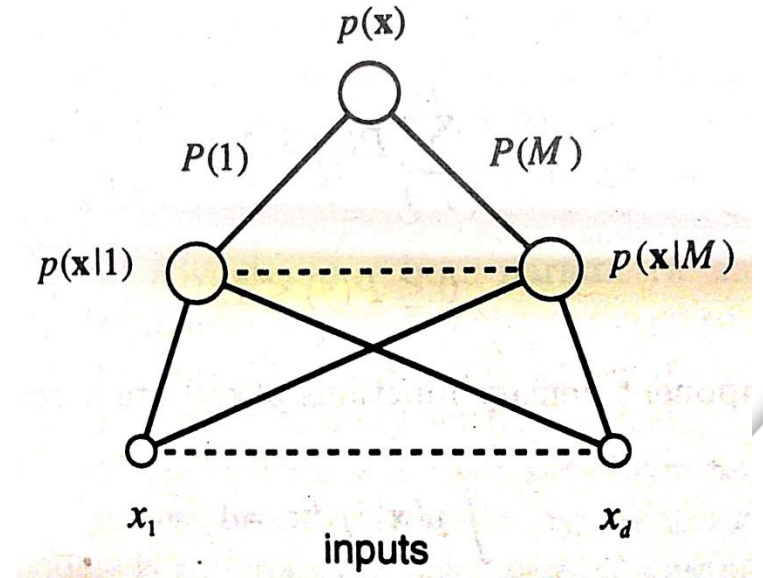
Incomplete Data : Component labels are not known

Classifier! $p(x|c_k)$ — class specific model

From Bayes' theorem:
$$P(\tau|x) = \frac{p(x|\tau) p(\tau)}{p(x)}$$

$$\sum_{\tau=1}^M p(\tau|x) = 1$$

$$p(x|\tau) = \frac{1}{(2\pi\sigma_\tau^2)^{\frac{d}{2}}} \exp \left\{ -\frac{\|x - \mu_\tau\|^2}{2\sigma_\tau^2} \right\}$$



Mixture Models : Maximum Likelihood Technique

Maximum Likelihood Technique

$$E = -\ln(L) = -\sum_{n=1}^N \ln p(x^n) = -\sum_{n=1}^N \ln \left[\sum_{\tau=1}^M p(x^n | \tau) p(\tau) \right]$$

Maximizing Likelihood $\Rightarrow$ Minimizing the Error f_m .

$\exists$ parameter $\Rightarrow$ Likelihood $(L) \rightarrow \infty$

$\mu = \mu_{\tau} ; \sigma_{\tau} = 0 \Rightarrow L \rightarrow \infty$

Singular Solutions

Local Minima : Small group of pts close together
poor rep of true distribution

Singular Solns Can be avoided by

- Constrain the components have equal co-variance
- Shrinking gaussian replaced by (angle over during iteration

Standard Non-linear optimization method for computing the minima - gradient info.

Mixture Models : ML Technique (Cont..)

ML Technique (cont..)

$$\frac{\partial E}{\partial \mu_j} = \sum_{n=1}^N P(j|x^n) \frac{(\mu_j - x^n)}{\sigma_j^2}$$

$$\frac{\partial E}{\partial \sigma_j^2} = \sum_{n=1}^N P(j|x^n) \left\{ \frac{d}{\sigma_j^2} - \frac{\|x^n - \mu_j\|^2}{\sigma_j^4} \right\}$$

$$P(j) = \frac{\exp(-\gamma_j)}{\sum_{k=1}^K \exp(-\gamma_k)} \quad -\infty \leq \gamma_j \leq \infty$$

$$\frac{\partial E}{\partial \gamma_j} = - \sum_{n=1}^N \{ P(j|x^n) - P(j) \}$$

$$\frac{\partial E}{\partial \mu_j} = 0 \Rightarrow \hat{\mu}_j = \frac{\sum_n P(j|x^n) x^n}{\sum_n P(j|x^n)}$$

$$\frac{\partial E}{\partial \sigma_j^2} = 0 \Rightarrow \hat{\sigma}_j^2 = \frac{1}{d} \left[\frac{\sum_n P(j|x^n) \|x^n - \hat{\mu}_j\|^2}{\sum_n P(j|x^n)} \right]$$

$$\frac{\partial E}{\partial \gamma_j} = 0 \Rightarrow \hat{P}_j = \frac{1}{N} \sum_n P(j|x^n)$$

Mixture Models : The EM Algorithm

The EM Algorithm

No direct soln to compute parameters $\hat{\mu}_T, \hat{\sigma}_T^2 \& \hat{P}(T)$

Highly non-linear coupled eqs $\Rightarrow$ Difficult to solve.

Iterative Scheme to find minimum of E

- Initiate with (guess) some parameters (known as "old")
- Evaluate RHS of $\hat{\mu}_T, \hat{\sigma}_T^2 \& \hat{P}(T) \Rightarrow$ New parameters
- New parameters $\Rightarrow$ Error $f_m \downarrow$
- New parameters $\rightarrow$ old parameters

Repeat above steps until local min is reached.

$$E_{\text{new}} - E^{\text{old}} = - \sum_n \ln \left[\frac{p^{\text{new}}(x^n)}{p^{\text{old}}(x^n)} \right]$$

$p^{\text{new}}(x) \& p^{\text{old}}(x) \Rightarrow$ Densities evaluated using new & old parameters respectively.

$$E^{\text{new}} - E^{\text{old}} = - \sum_n \ln \left[\frac{\sum_T p^{\text{new}}(T) p^{\text{new}}(x^n|T)}{p^{\text{old}}(x^n)} \cdot \frac{p^{\text{old}}(T|x^n)}{p^{\text{old}}(T|x^n)} \right]$$

Mixture Models : The EM Algorithm (Cont..)

The EM Algorithm (cont..)

using Jensen's inequality $\lambda_J \geq 0$ & $\sum_J \lambda_J = 1$

$$\ln \left(\sum_J \lambda_J x_J \right) \geq \sum_J \lambda_J \ln(x_J)$$

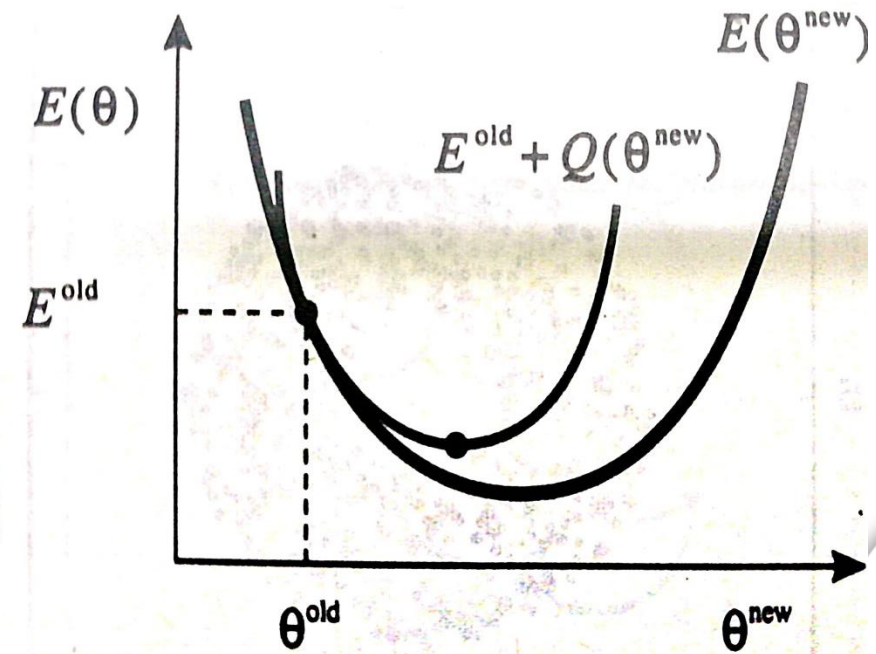
Applying JE on $E^{\text{new}} - E^{\text{old}}$

$$E^{\text{new}} - E^{\text{old}} \leq - \sum_J \sum_m p^{\text{old}}(J|x^m) \ln \left[\frac{p^{\text{new}}(J) p^{\text{new}}(x^m|J)}{p^{\text{old}}(x^m) p^{\text{old}}(J|x^m)} \right]$$

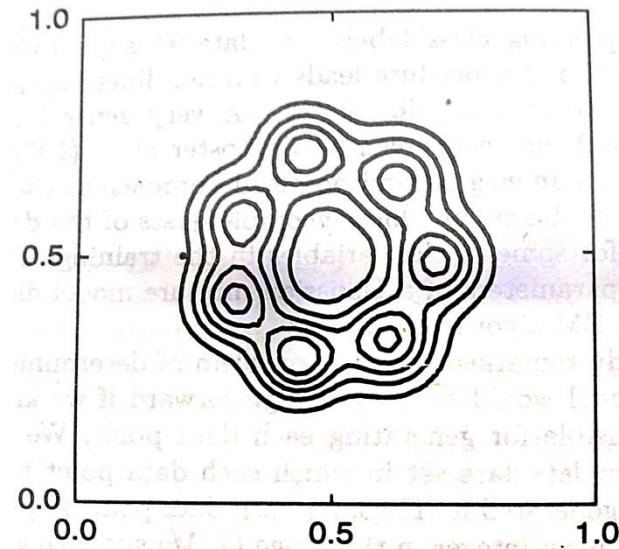
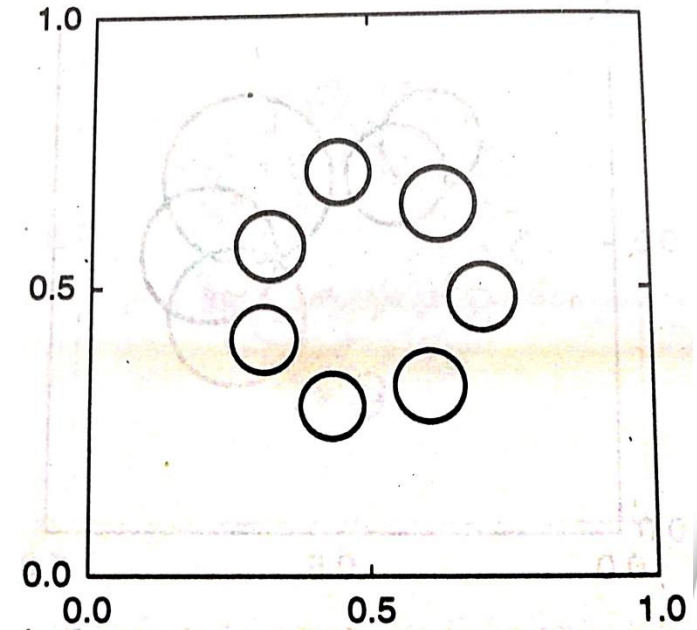
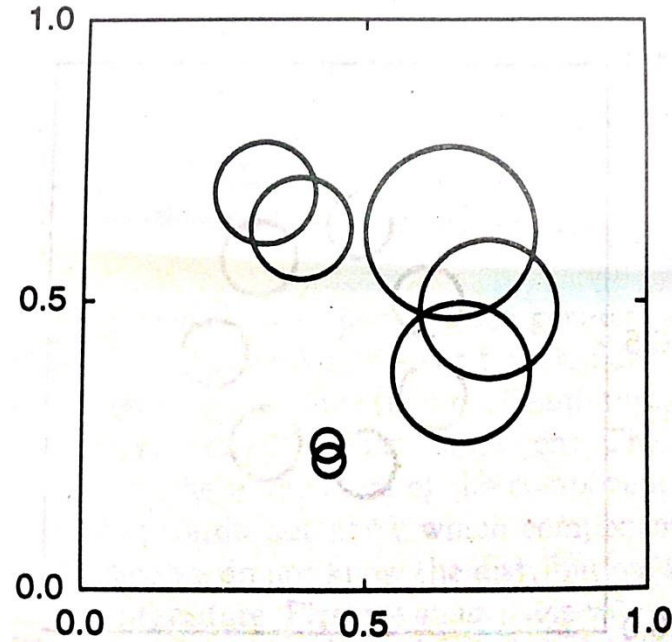
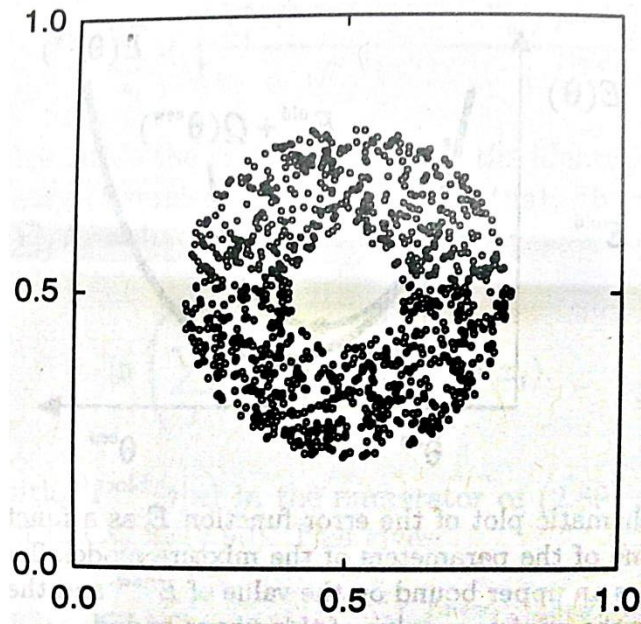
$$\mu_J^{\text{new}} = \frac{\sum_m p^{\text{old}}(J|x^m) x^m}{\sum_m p^{\text{old}}(J|x^m)}$$

$$(\sigma_J^{\text{new}})^2 = \frac{1}{2} \left[\frac{\sum_m p^{\text{old}}(J|x^m) \|x^m - \mu_J^{\text{new}}\|^2}{\sum_m p^{\text{old}}(J|x^m)} \right]$$

$$P_J^{\text{new}} = \frac{1}{2} \sum_m p^{\text{old}}(J|x^m)$$



Mixture Models : The EM Algorithm (Example)



Mixture Models : The EM Algorithm (Incomplete Data)

EM Algorithm for Incomplete Data

Hypothetical ^{complete} Dataset $\Rightarrow [x^n : z^n]$; $n = 1, 2, \dots, N$
 $z^n = 1, 2, \dots, M$

$z^n \rightarrow$ Label of the Mixture component (1, M)

$$E^{\text{comp}} = -\ln L^{\text{comp}} = -\sum_{n=1}^N \ln p^{\text{new}}(x^n, z^n)$$

$$= -\sum_{n=1}^N \ln \left[p^{\text{new}}(z^n) p^{\text{new}}(x^n | z^n) \right]$$

$\therefore$ we don't have the PD of z^n , we estimate by EM algo

- Given some parameters of Mixture Model (old values)
- use the above parameters & Bayes' theorem find PD of z^n
- Then compute Expectation of E^{comp}

$$E\text{-Step} \Rightarrow E[E^{\text{comp}}]$$

- New values of parameters are determined by minimizing $E[E^{\text{comp}}]$ w.r.t parameters

Mixture Models : The EM Algorithm (Incomplete Data)

EM Algorithm - Incomplete Data (cont.)

Minimization of $E[E^{\text{comp}}] \Rightarrow$ Maximization of $E[L^{\text{comp}}]$

M-Step $\Rightarrow$ Maximization of $E[L^{\text{comp}}]$

$$\begin{aligned} E[E^{\text{comp}}] &= \sum_{z^1=1}^M \sum_{z^2=1}^M \dots \sum_{z^N=1}^M E^{\text{comp}} \prod_{n=1}^N p^{\text{old}}(z^n | x^n) \\ &= - \sum_{n=1}^N \sum_{j=1}^M p^{\text{old}}(j | x^n) \ln \left[p^{\text{new}}(j) p^{\text{new}}(x^n | j) \right] \end{aligned}$$

Illustration of Use of GMMs : Isolated Word Rec

Consider GMMs
In place of HMMs

