



Statistics: The Backbone of Statistical Learning

Alexander G. Ororbia II and Viet Nguyen
Introduction to Machine Learning
CSCI-335
9/16/2026

Statistics

Standard deviation - a measure of data points differ from mean

Average of differences (w/ mean as reference point)

Higher standard deviation indicates higher spread, less consistency, and less “clustering/blobbing”

Sample standard deviation: $s = \sqrt{\frac{\sum (x - \bar{X})^2}{n - 1}}$ **Mean (average):** $\bar{X} = \left(\frac{1}{n}\right) \sum (x)$

Population standard deviation: $\sigma = \sqrt{\frac{\sum (x - \mu)^2}{N}}$

Expectation

X is univariate here!

$$\mathbb{E}[X] = \sum_{i=1}^k x_i p_i = x_1 p_1 + x_2 p_2 + \cdots + x_k p_k.$$

Definition A random vector $\vec{X}$ is a vector $(X_1, X_2, \dots, X_p)$ of jointly distributed random variables. As is customary in linear algebra, we will write vectors as column matrices whenever convenient.

Definition The expectation $E\vec{X}$ of a random vector $\vec{X} = [X_1, X_2, \dots, X_p]^T$ is given by

$$E\vec{X} = \begin{bmatrix} EX_1 \\ EX_2 \\ \vdots \\ EX_p \end{bmatrix}.$$

The linearity properties of the expectation can be expressed compactly by stating that for any $k \times p$ -matrix A and any $1 \times j$ -matrix B ,

$$E(A\vec{X}) = AE\vec{X} \quad \text{and} \quad E(\vec{X}B) = (E\vec{X})B.$$

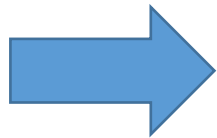
X is $p \times 1$ here!

X is $j \times 1$ here!

$$\begin{aligned} E\vec{X} &= E \begin{bmatrix} X_1 \\ 0 \\ \vdots \\ 0 \end{bmatrix} + E \begin{bmatrix} 0 \\ X_2 \\ \vdots \\ 0 \end{bmatrix} + \cdots + E \begin{bmatrix} 0 \\ 0 \\ \vdots \\ X_p \end{bmatrix} \\ &= \begin{bmatrix} EX_1 \\ 0 \\ \vdots \\ 0 \end{bmatrix} + \begin{bmatrix} 0 \\ EX_2 \\ \vdots \\ 0 \end{bmatrix} + \cdots + \begin{bmatrix} 0 \\ 0 \\ \vdots \\ EX_p \end{bmatrix} \\ &= \begin{bmatrix} EX_1 \\ EX_2 \\ \vdots \\ EX_p \end{bmatrix}. \end{aligned}$$

Covariance

- Variance and Covariance:
 - Measure of the “spread” of a set of points around their center of mass (mean)
- Variance:
 - Measure of deviation from mean for points in one dimension
- Covariance:
 - Measure of how much each of dimensions vary from mean with **respect to each other**



- **Covariance is measured between two dimensions**
- **Covariance → relation between two dimensions**
- **Covariance between one dimension is variance**

Variance-Covariance

Definition The *variance-covariance matrix* (or simply the *covariance matrix*) of a random vector $\vec{X}$ is given by:

$$\text{Cov}(\vec{X}) = E \left[(\vec{X} - E\vec{X})(\vec{X} - E\vec{X})^T \right].$$

Proposition

$$\text{Cov}(\vec{X}) = E[\vec{X}\vec{X}^T] - E\vec{X}(E\vec{X})^T.$$

Proposition

$$\text{Cov}(\vec{X}) = \begin{bmatrix} \text{Var}(X_1) & \text{Cov}(X_1, X_2) & \cdots & \text{Cov}(X_1, X_p) \\ \text{Cov}(X_2, X_1) & \text{Var}(X_2) & \cdots & \text{Cov}(X_2, X_p) \\ \vdots & \vdots & \ddots & \vdots \\ \text{Cov}(X_p, X_1) & \text{Cov}(X_p, X_2) & \cdots & \text{Var}(X_p) \end{bmatrix}.$$

Thus, $\text{Cov}(\vec{X})$ is a symmetric matrix, since $\text{Cov}(X, Y) = \text{Cov}(Y, X)$.

Covariance Matrix

- Gives a measure of the dispersion of the data
- It is a $D \times D$ matrix
 - Element in position i,j is the covariance between the i^{th} and j^{th} variables.
 - Covariance between two variables x_i and x_j is defined as $E[(x_i - \mu_i)(x_j - \mu_j)]$
 - Can be positive or negative
 - If the variables are independent then the covariance is zero.
 - Then all matrix elements are zero except diagonal elements which represent the variances

The Gaussian Distribution



Carl Friedrich Gauss
1777-1855

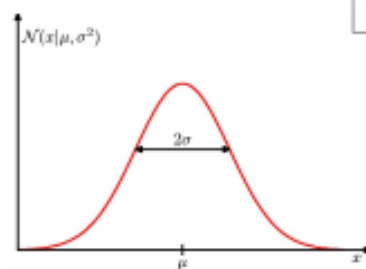
- For single real-valued variable x

$$N(x | \mu, \sigma^2) = \frac{1}{(2\pi\sigma^2)^{1/2}} \exp\left\{-\frac{1}{2\sigma^2}(x - \mu)^2\right\}$$

- Parameters:

- Mean μ , variance σ^2 ,

- *Standard deviation* σ



68% of data lies within σ of mean
95% within 2σ

- *Precision* $\beta = 1/\sigma^2$, $E[x] = \mu$, $Var[x] = \sigma^2$

- For D -dimensional vector $\mathbf{x}$, multivariate Gaussian

$$N(\mathbf{x} | \boldsymbol{\mu}, \Sigma) = \frac{1}{(2\pi)^{D/2}} \frac{1}{|\Sigma|^{1/2}} \exp\left\{-\frac{1}{2}(\mathbf{x} - \boldsymbol{\mu})^T \Sigma^{-1}(\mathbf{x} - \boldsymbol{\mu})\right\}$$

$\boldsymbol{\mu}$ is a mean vector, Σ is a $D \times D$ covariance matrix, $|\Sigma|$ is the determinant of Σ

Σ^{-1} is also referred to as the precision matrix

Matrix Determinant

$$A = \begin{bmatrix} a & b & c \\ d & e & f \\ g & h & i \end{bmatrix}$$

$$|A| = a(ei - fh) - b(di - fg) + c(dh - eg)$$

"The determinant of A equals ... etc"

$$\begin{bmatrix} a_x & e & f \\ d & b_x & f \\ g & h & i \end{bmatrix} - \begin{bmatrix} d & b_x & f \\ g & h & i \end{bmatrix} + \begin{bmatrix} d & e & x_c \\ g & h & i \end{bmatrix}$$

The Gaussian Distribution



Carl Friedrich Gauss
1777-1855

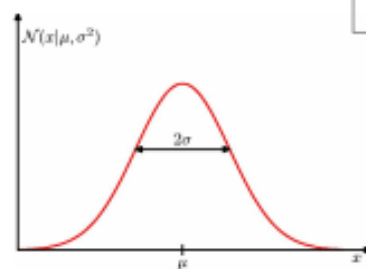
- For single real-valued variable x

$$N(x | \mu, \sigma^2) = \frac{1}{(2\pi\sigma^2)^{1/2}} \exp\left\{-\frac{1}{2\sigma^2}(x - \mu)^2\right\}$$

- Parameters:

- Mean μ , variance σ^2 ,

- *Standard deviation* σ



68% of data lies within σ of mean
95% within 2σ

- *Precision* $\beta = 1/\sigma^2$, $E[x] = \mu$, $Var[x] = \sigma^2$

- For D -dimensional vector $\mathbf{x}$, multivariate Gaussian

$$N(\mathbf{x} | \boldsymbol{\mu}, \Sigma) = \frac{1}{(2\pi)^{D/2}} \frac{1}{|\Sigma|^{1/2}} \exp\left\{-\frac{1}{2}(\mathbf{x} - \boldsymbol{\mu})^T \Sigma^{-1}(\mathbf{x} - \boldsymbol{\mu})\right\}$$

$\boldsymbol{\mu}$ is a mean vector, Σ is a $D \times D$ covariance matrix, $|\Sigma|$ is the determinant of Σ

Σ^{-1} is also referred to as the precision matrix

The Multivariate Gaussian Distribution

A p -dimensional random vector $\vec{X}$ has the *multivariate normal distribution* if it has the density function

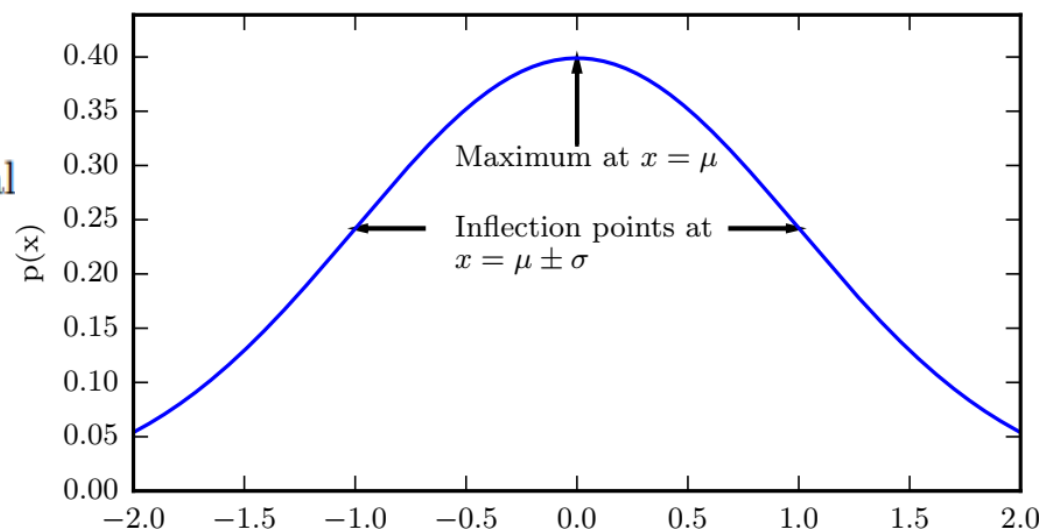
$$f(\vec{X}) = (2\pi)^{-p/2} |\Sigma|^{-1/2} \exp \left(-\frac{1}{2} (\vec{X} - \vec{\mu})^T \Sigma^{-1} (\vec{X} - \vec{\mu}) \right),$$

where $\vec{\mu}$ is a constant vector of dimension p and Σ is a $p \times p$ positive semi-definite which is invertible (called, in this case, *positive definite*). Then, $E\vec{X} = \vec{\mu}$ and $\text{Cov}(\vec{X}) = \Sigma$.

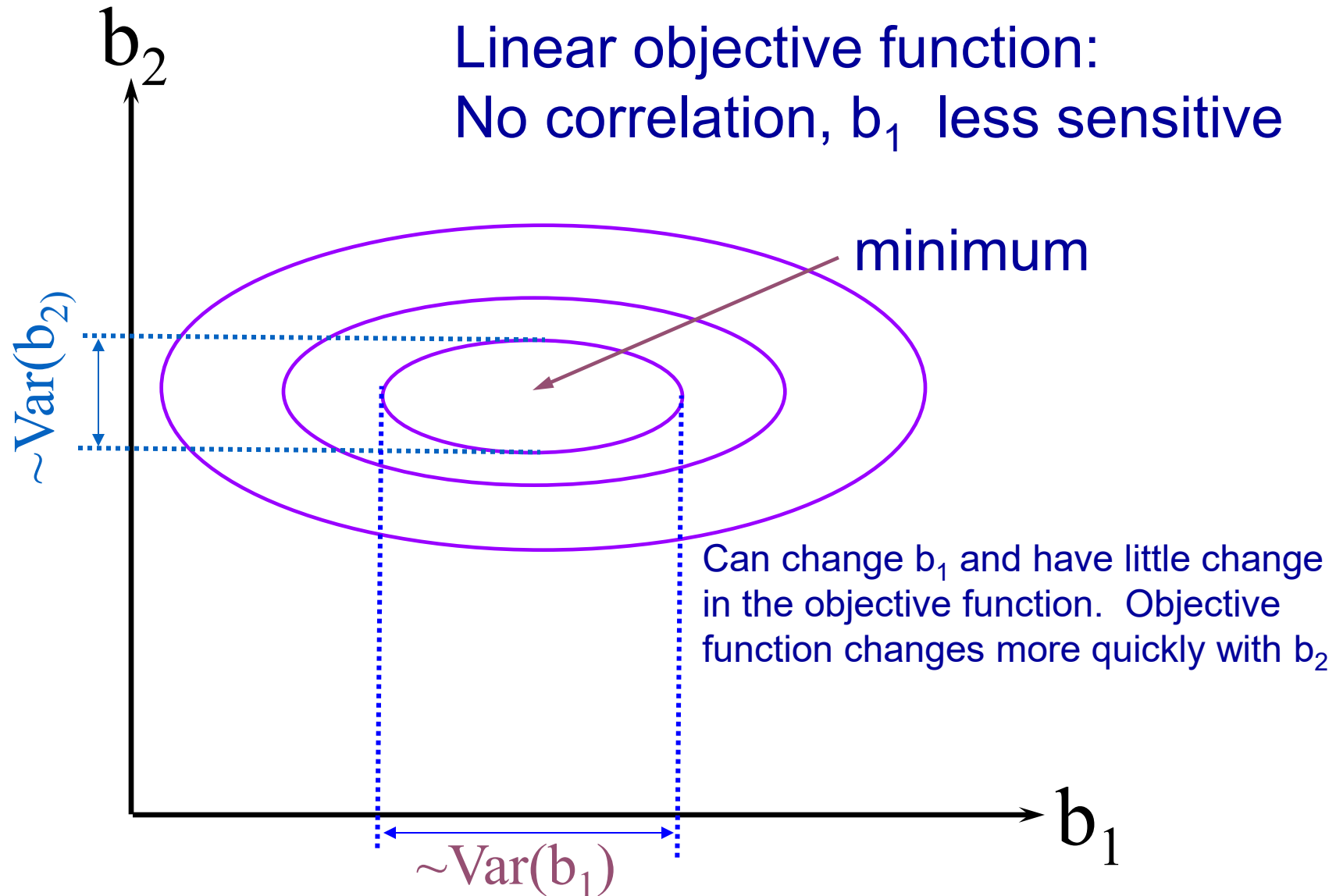
The *standard multivariate normal distribution* is obtained when $\vec{\mu} = 0$ and $\Sigma = I_p$, the $p \times p$ identity matrix:

$$f(\vec{X}) = (2\pi)^{-p/2} \exp \left(-\frac{1}{2} \vec{X}^T \vec{X} \right).$$

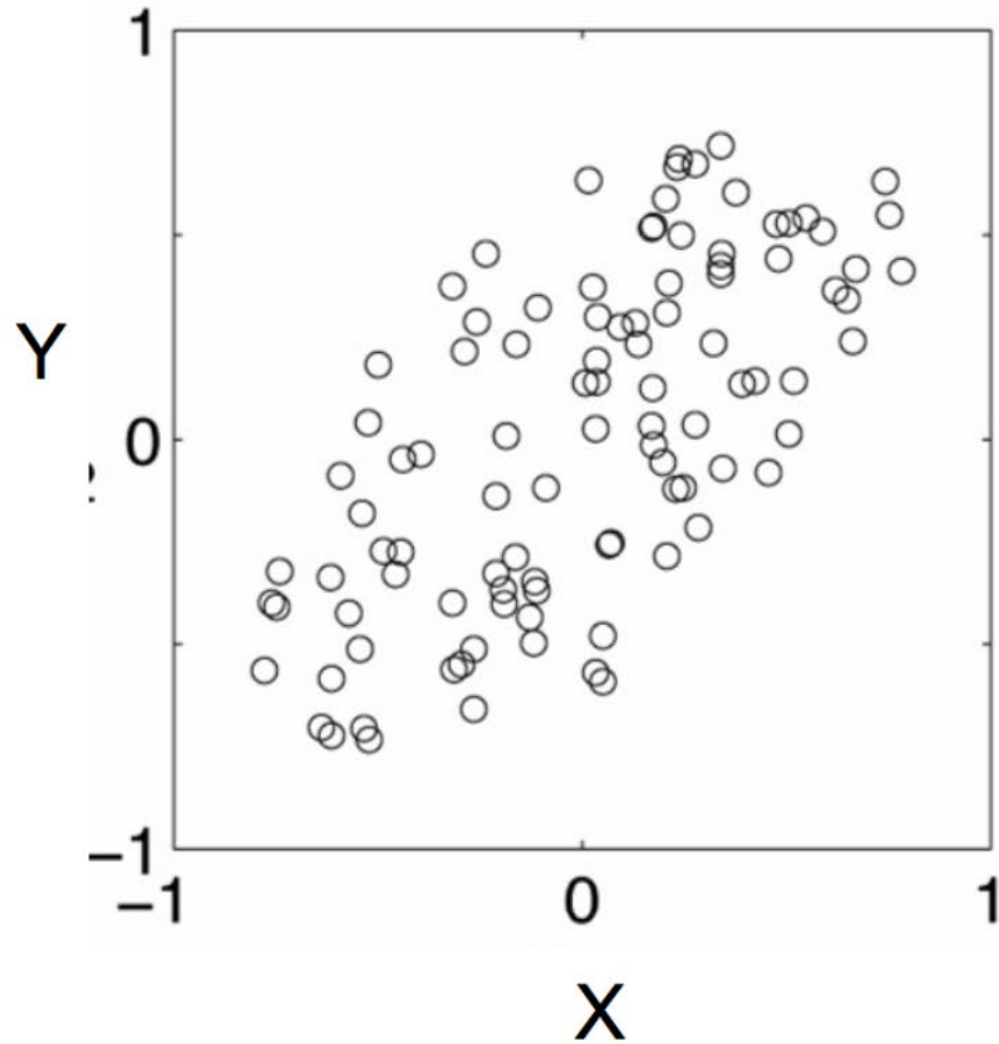
This corresponds to the case where $X_1, \dots, X_p$ are i.i.d. standard normal



Example: Parameter Variance & Covariance

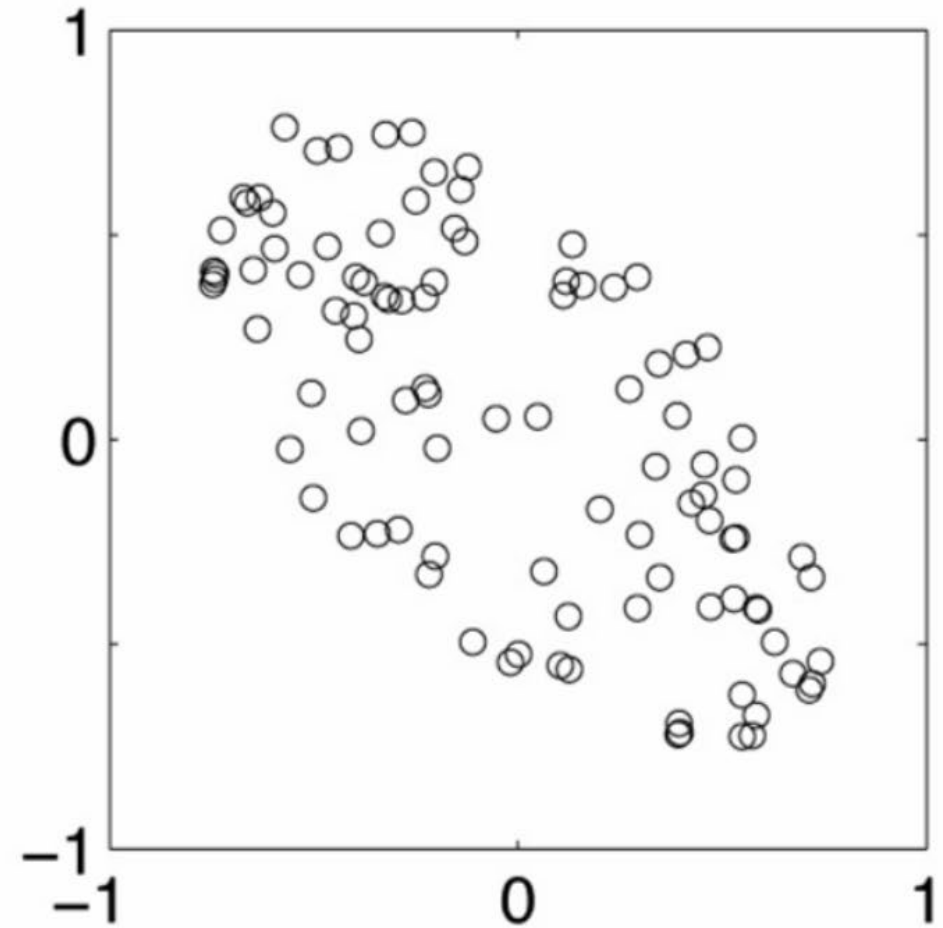


positive covariance



Positive: Both dimensions increase or decrease together

negative covariance



Negative: While one increases the other decreases

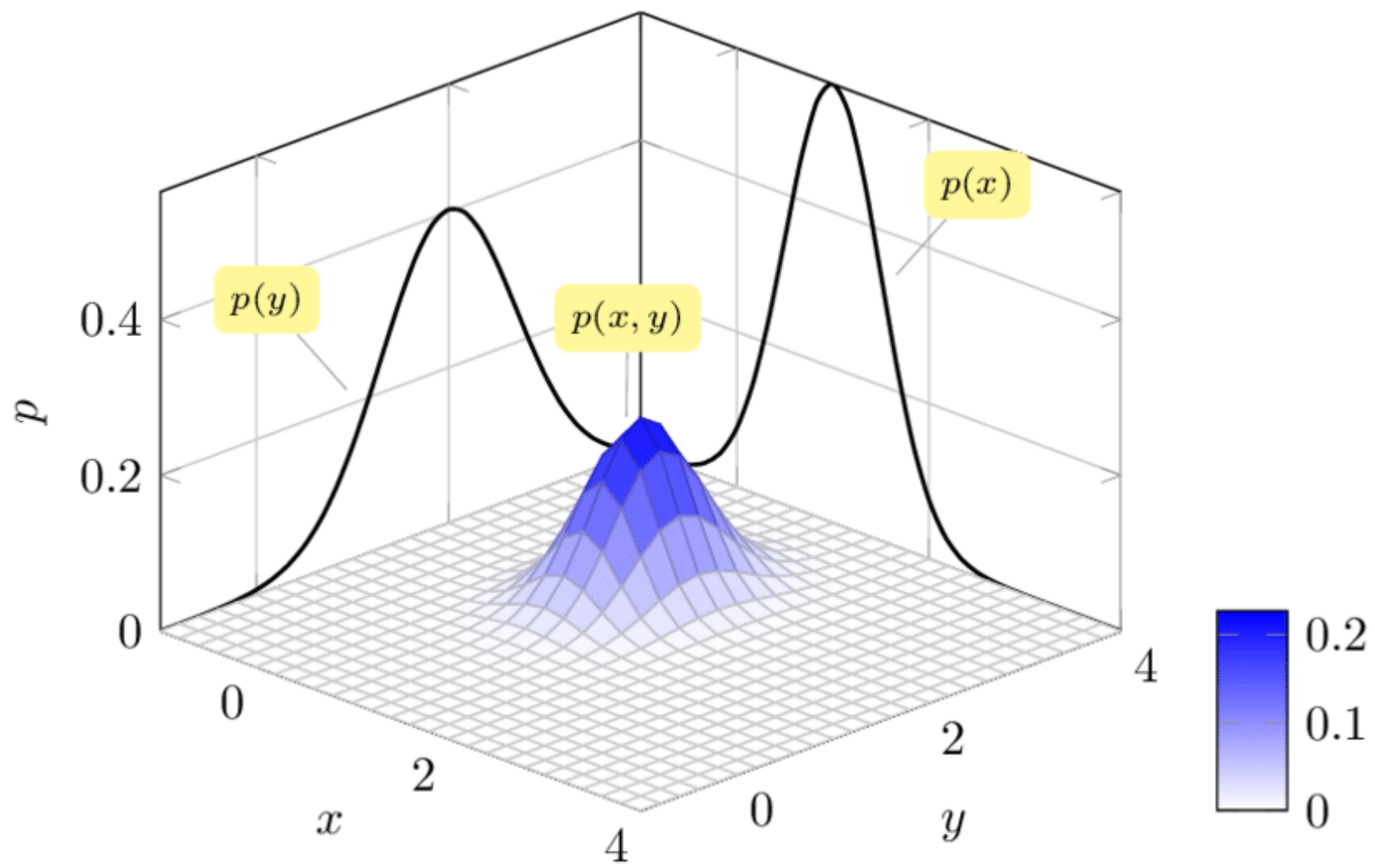
Covariance

- Used to find relationships between dimensions in high dimensional data sets

$$q_{jk} = \frac{1}{N} \sum_{i=1}^N (X_{ij} - E(X_j)) (X_{ik} - E(X_k))$$



The sample mean



Anomaly detection with the multivariate Gaussian

1. Fit model $p(x)$ by setting

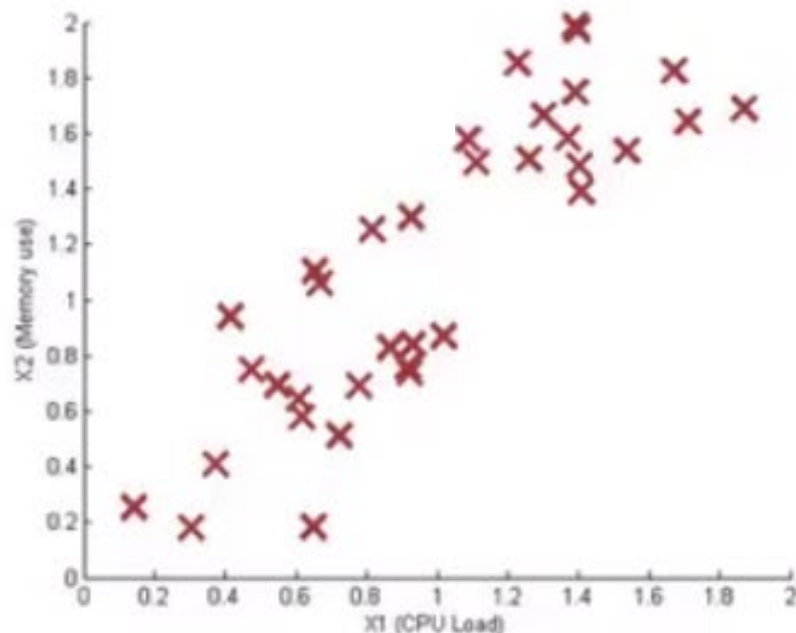
$$\mu = \frac{1}{m} \sum_{i=1}^m x^{(i)}$$

$$\Sigma = \frac{1}{m} \sum_{i=1}^m (x^{(i)} - \mu)(x^{(i)} - \mu)^T$$

2. Given a new example x , compute

$$p(x) = \frac{1}{(2\pi)^{\frac{n}{2}} |\Sigma|^{\frac{1}{2}}} \exp \left(-\frac{1}{2} (x - \mu)^T \Sigma^{-1} (x - \mu) \right)$$

Flag an anomaly if $p(x) < \varepsilon$

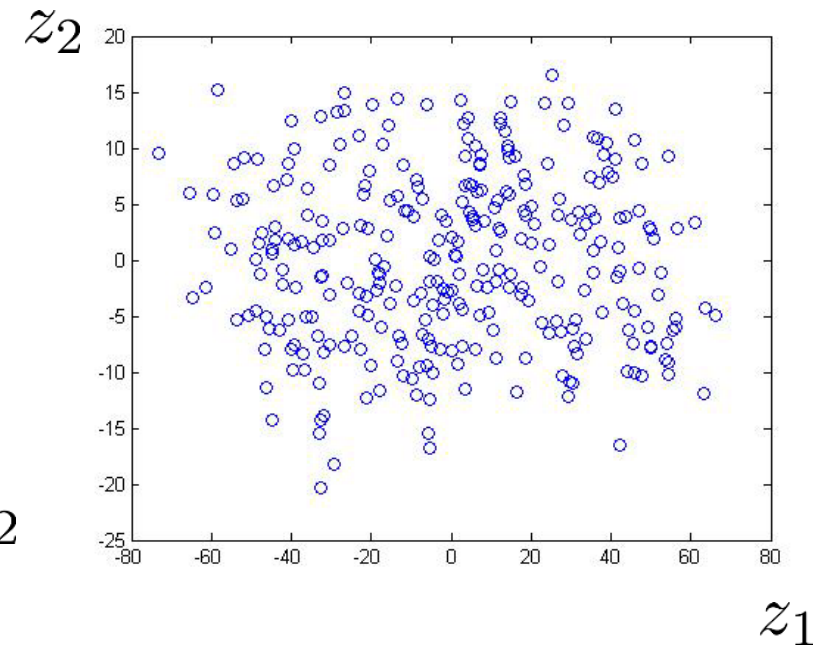
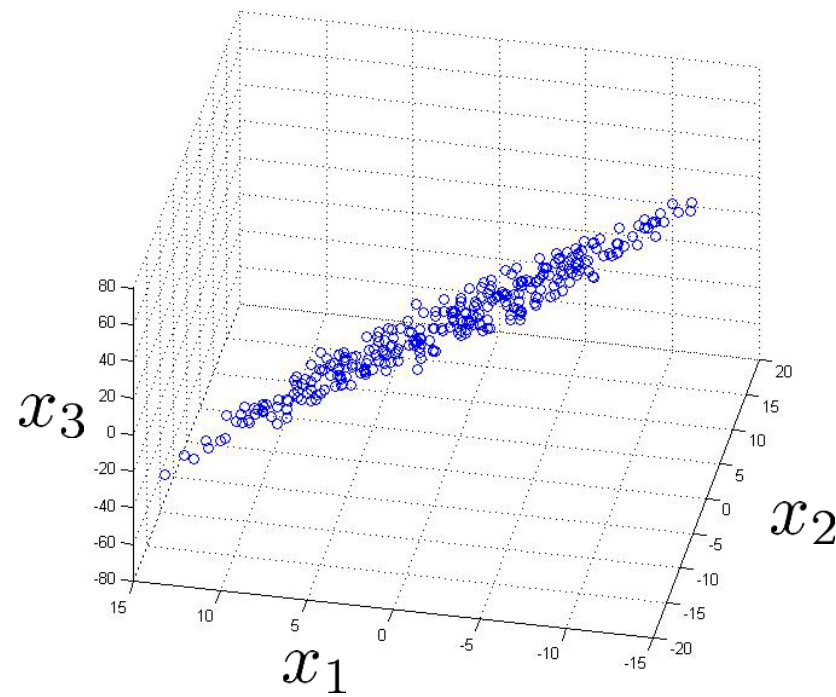
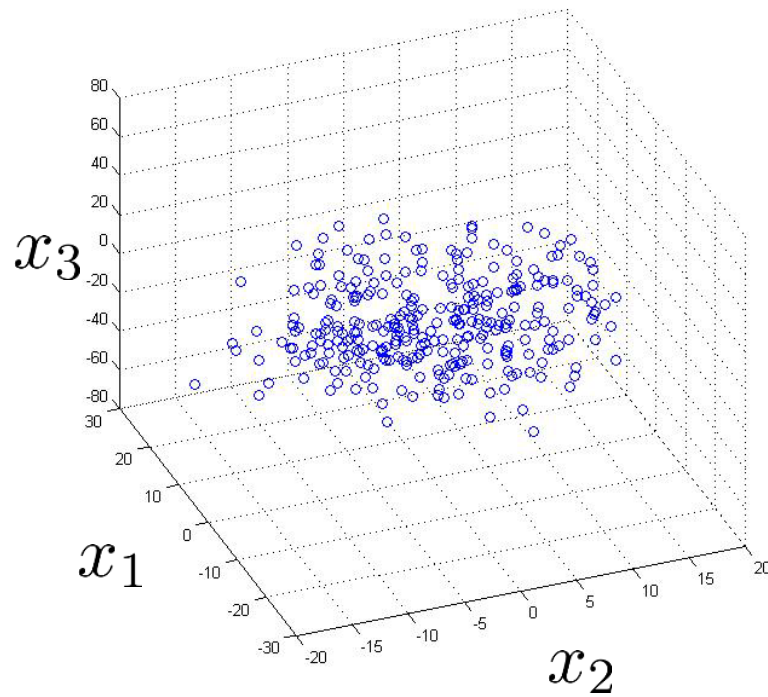


Why? Well, Dimensionality Reduction...

- PCA (Principal Component Analysis):
 - Find projection that maximize the variance (based on Gaussian assumptions)
- ICA (Independent Component Analysis):
 - Similar to PCA except assumes non-Gaussian features
- Multidimensional Scaling:
 - Find projection that best preserves inter-point distances
- LDA (Linear Discriminant Analysis):
 - Maximizing the component axes for class-separation
- ...

Data Compression

Reduce data from 3D to 2D

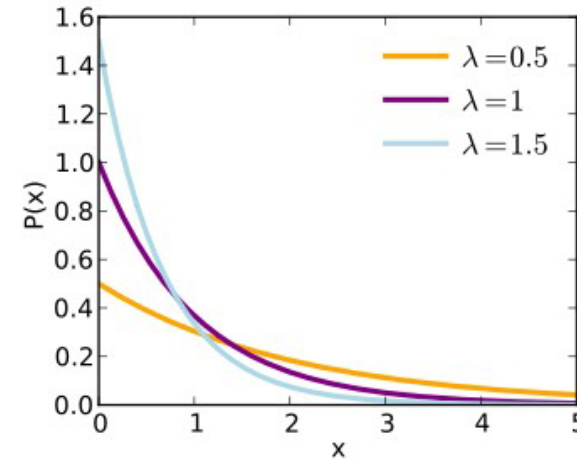


More Distributions

Exponential:

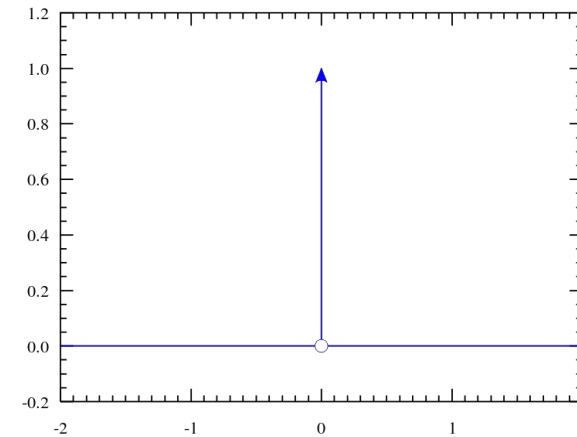
$$p(x; \lambda) = \lambda \mathbf{1}_{x \geq 0} \exp(-\lambda x).$$

Used to predict the waiting time until the next event occurs, such as a success, failure, or arrival



Dirac

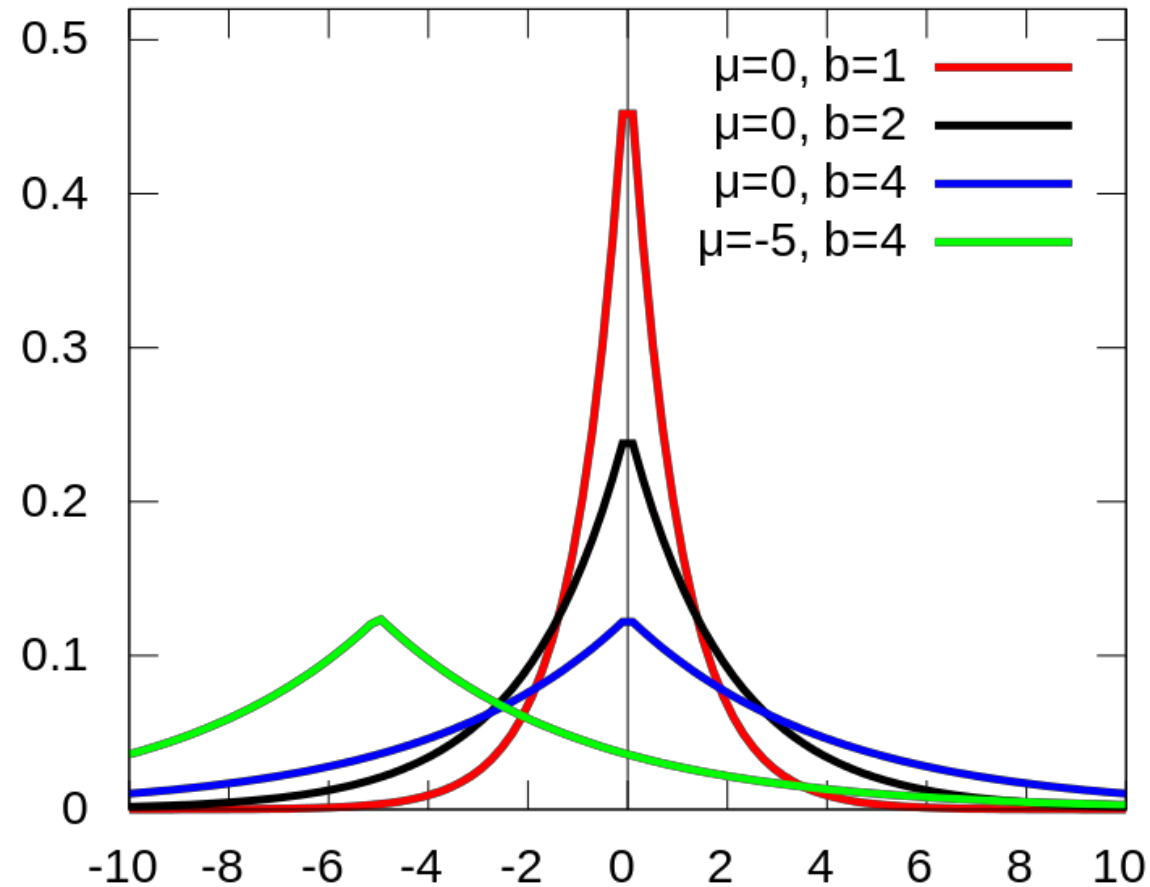
$$p(x) = \delta(x - \mu)$$



“Dirac density” of an idealized point mass or point charge -- a function that is equal to zero everywhere except for zero (integral over the entire real line is equal to one)

Laplace Distribution

$$\text{Laplace}(x; \mu, \gamma) = \frac{1}{2\gamma} \exp\left(-\frac{|x - \mu|}{\gamma}\right)$$



Bernoulli Distribution

The Bernoulli distribution is the “coin flip” distribution.

X is Bernoulli if its probability function is:

$$X = \begin{cases} 1 & \text{w.p. } p \\ 0 & \text{w.p. } 1-p \end{cases}$$

$X=1$ is usually interpreted as a “success.” E.g.:

$X=1$ for heads in coin toss

$X=1$ for male in survey

$X=1$ for defective in a test of product

$X=1$ for “made the sale” tracking performance

Bernoulli Distribution

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

$$P(x = 0) = 1 - \phi$$

$$P(x = x) = \phi^x (1 - \phi)^{1-x}$$

$$\mathbb{E}_x[x] = \phi$$

$$\text{Var}_x(x) = \phi(1 - \phi)$$

Can prove/derive each of these properties!

p is ϕ in these formulas!

Empirical Distribution

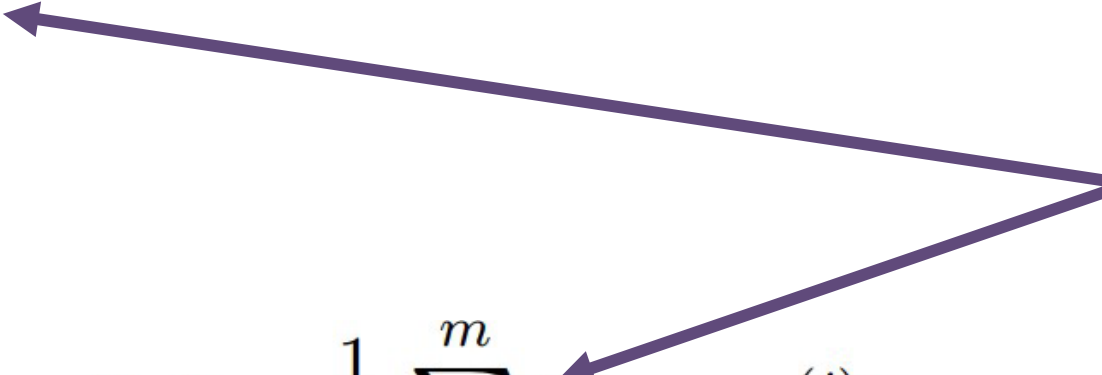
$$\delta(x) = \begin{cases} 0, & x \neq 0 \\ \infty, & x = 0 \end{cases}$$

such that:

$$\int_{-\infty}^{\infty} \delta(x) dx = 1$$

$$\hat{p}(\mathbf{x}) = \frac{1}{m} \sum_{i=1}^m \delta(\mathbf{x} - \mathbf{x}^{(i)})$$

Dirac delta function



An **empirical (Dirac) distribution function** is
distribution function associated with empirical
measure of a sample
(the data “is” the distribution)

Mixture Distributions

$$P(\mathbf{x}) = \sum_i P(c = i) P(\mathbf{x} \mid c = i)$$

Gaussian mixture with
three components

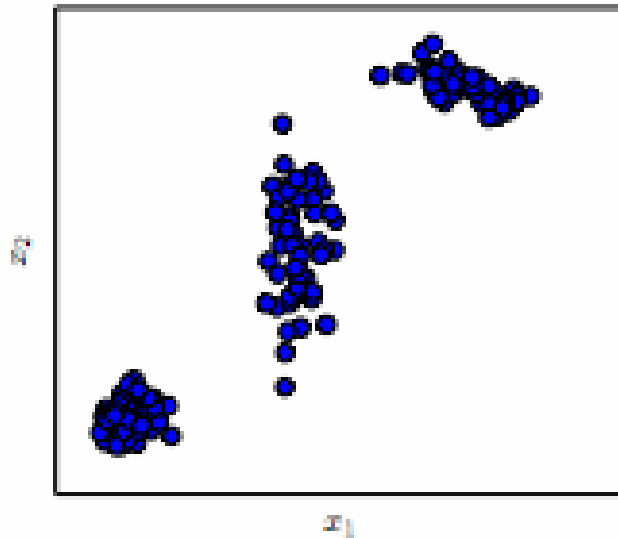


Figure 3.2

Mixtures of Gaussians

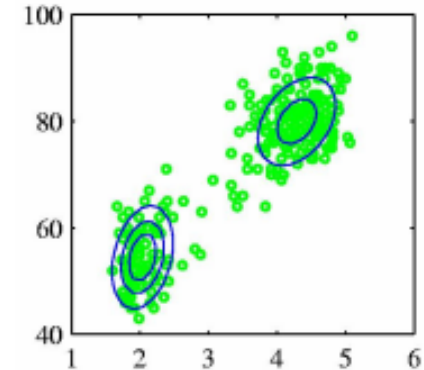
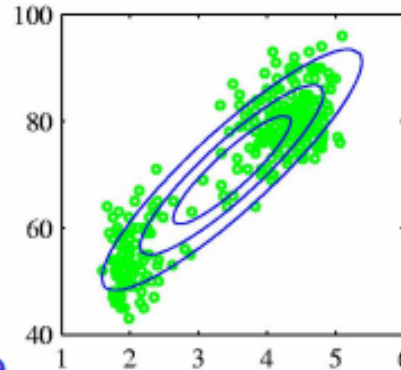
- Gaussian has limitations in modeling real data sets
- Old Faithful (Hydrothermal Geyser in Yellowstone)
 - 272 observations
 - Duration (mins, horiz axis) vs Time to next eruption (vertical axis)
 - Simple Gaussian unable to capture structure
 - Linear superposition of two Gaussians is better



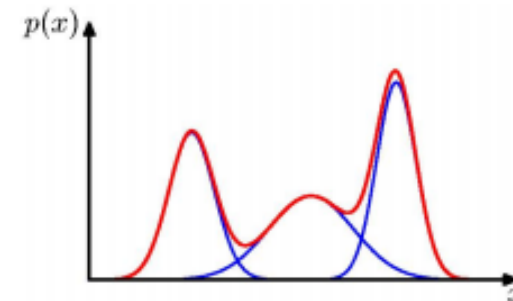
- Linear combinations of Gaussians can give very complex densities

$$p(x) = \sum_{k=1}^K \pi_k N(x | \mu_k, \Sigma_k)$$

π_k are mixing coefficients that sum to one



- One –dimension
 - Three Gaussians in blue
 - Sum in red



We will build this model later in this class!

Questions?

Deep robots!

Deep questions?!

