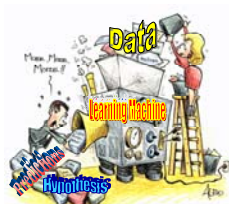


Advanced Machine Learning

Introduction

Eric Xing

Lecture 1, August 10, 2009



Reading:

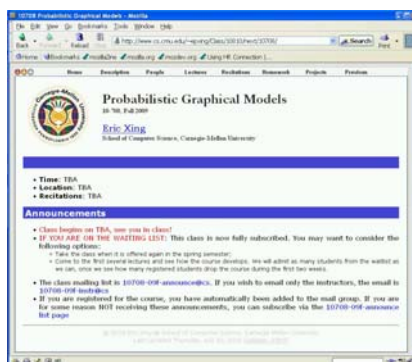
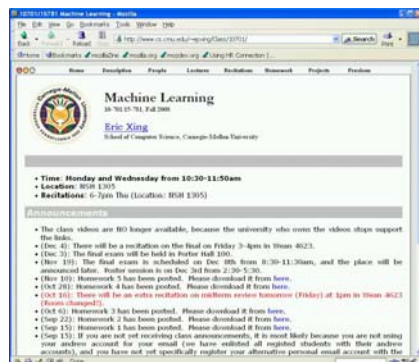
© Eric Xing @ CMU, 2006-2009



Advanced Machine Learning

- Where does it come from:

- <http://www.cs.cmu.edu/~epxing/Class/10701/>
- <http://www.cs.cmu.edu/~epxing/Class/10708/>



© Eric Xing @ CMU, 2006-2009



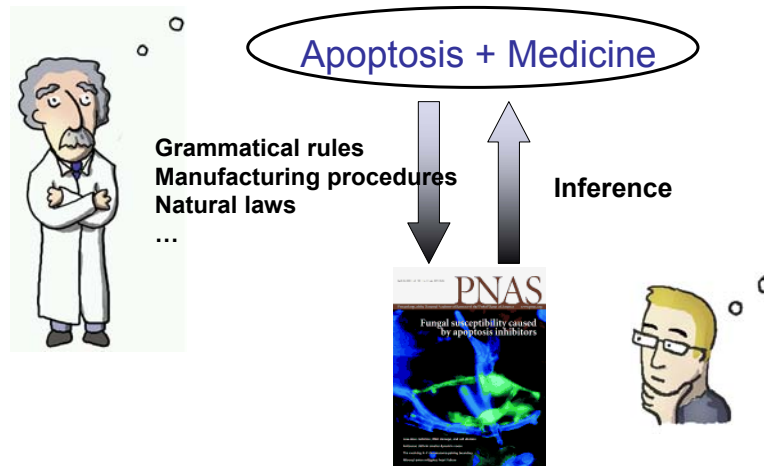
Logistics

- Text book
 - Chris Bishop, **Pattern Recognition and Machine Learning** (required)
 - Tom Mitchell, **Machine Learning**
 - David Mackay, **Information Theory, Inference, and Learning Algorithms**
 - Daphnie Koller and Nir Friedman, **Probabilistic Graphical Models**
- Class resource
 - <http://www.cis.pku.edu.cn/faculty/vision/wangliwei/DS.html>
 - <http://www.cs.cmu.edu/~epxing/Class/DS>
- TA:
 - Dr. Jun Zhu
 - ...
- Host:
 - Liwei Wang, Peking University
 - Changshui Zhang, Tsinghua University

© Eric Xing @ CMU, 2006-2009

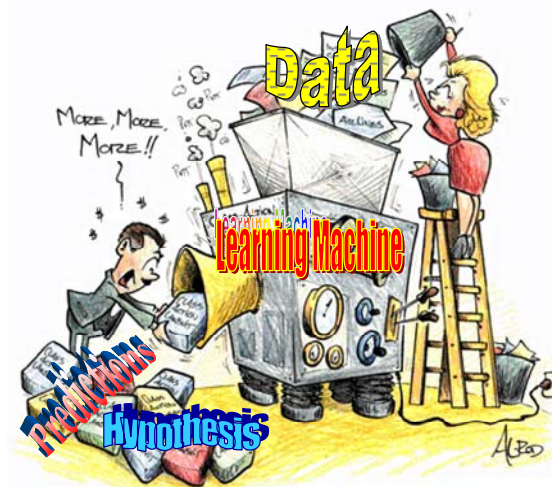
What is Learning

Learning is about seeking a **predictive** and/or **executable** understanding of natural/artificial subjects, phenomena, or activities from ...



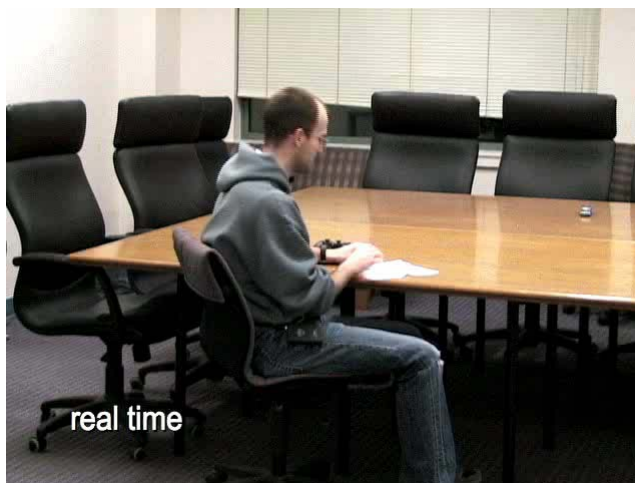
© Eric Xing @ CMU, 2006-2009

Machine Learning



© Eric Xing @ CMU, 2006-2009

Fetching a stapler from inside an office --- the Stanford STAIR robot



© Eric Xing @ CMU, 2006-2009

Machine Learning



Machine Learning seeks to develop theories and computer systems for

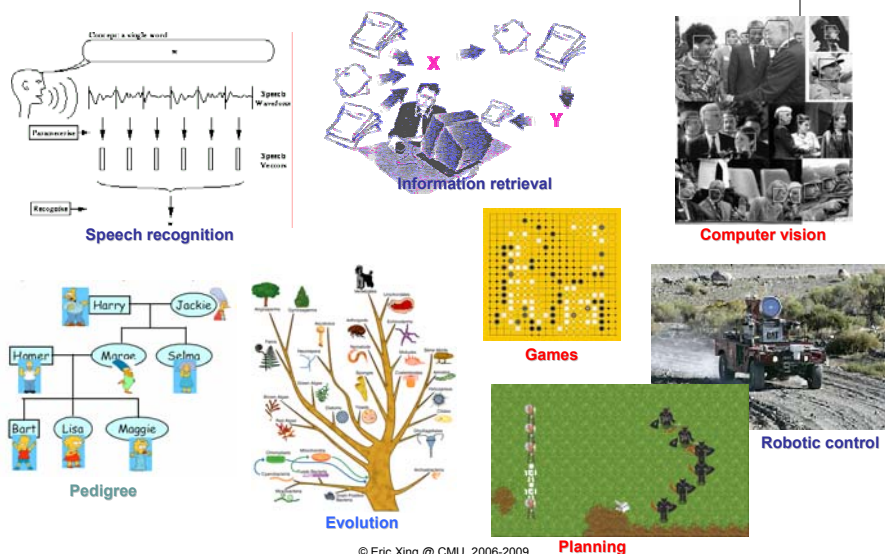
- representing;
- classifying, clustering and recognizing;
- reasoning under uncertainty;
- predicting;
- and reacting to
- ...

complex, real world data, based on the system's own experience with data, and (hopefully) under a unified model or mathematical framework, that

- can be formally characterized and analyzed
- can take into account human prior knowledge
- can generalize and adapt across data and domains
- can operate automatically and autonomously
- and can be interpreted and perceived by human.

© Eric Xing @ CMU, 2006-2009

Where Machine Learning is being used or can be useful?



© Eric Xing @ CMU, 2006-2009

Natural language processing and speech recognition



- Now most pocket **Speech Recognizers** or **Translators** are running on some sort of learning device --- the more you play/use them, the smarter they become!

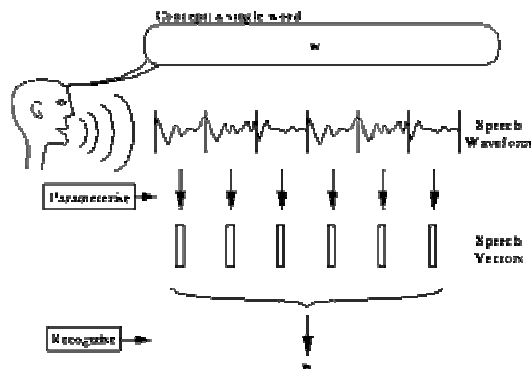


Fig. 1.2 Isolated Word Problem

© Eric Xing @ CMU, 2006-2009

Object Recognition



- Behind a security camera, most likely there is a computer that is learning and/or checking!



© Eric Xing @ CMU, 2006-2009

Robotic Control I



- The **best** helicopter pilot is now a computer!
 - it runs a program that learns how to fly and make acrobatic maneuvers by itself!
 - no taped instructions, joysticks, or things like ...



A. Ng 2005

© Eric Xing @ CMU, 2006-2009

Robotic Control II

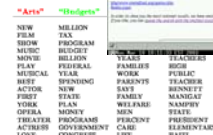


- Now cars can find their own ways!



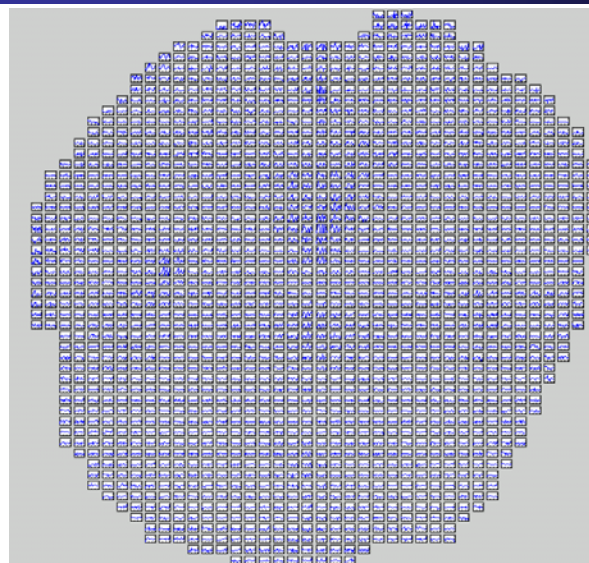
© Eric Xing @ CMU, 2006-2009

- **We want:**



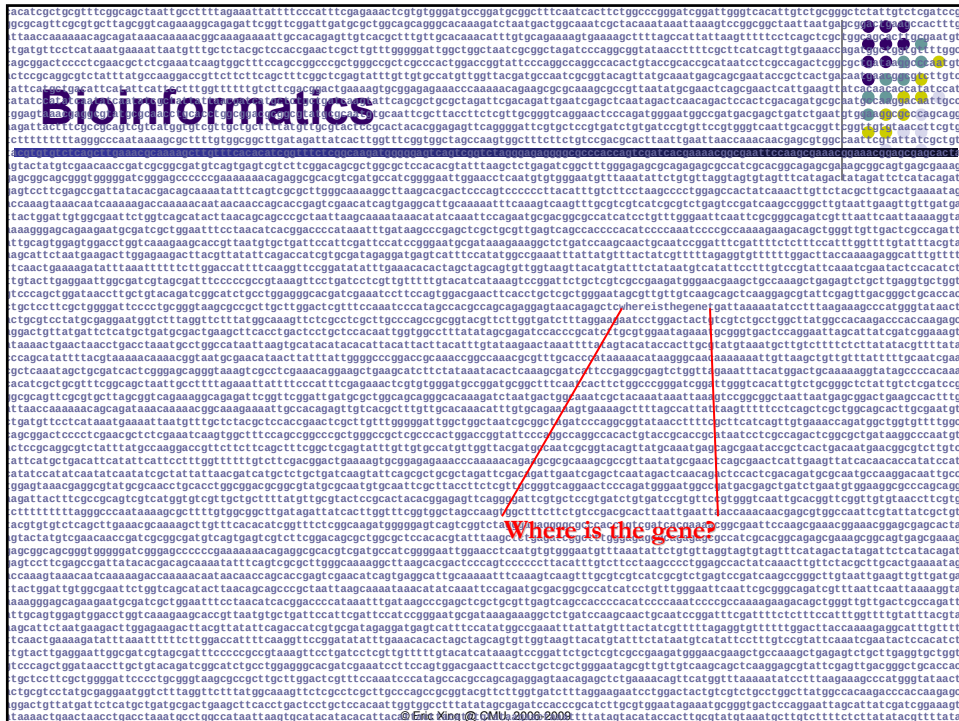
The Wilson-Randolph-Hoover Foundation will give **\$1.26 million** to **Lincoln Center**, **Metropolitan Opera Co.**, **New York Philharmonic** and **Juilliard School**. "Our heart felt wish that we had a net opportunity to make a mark on the future of the performing arts with these **great** **art** **sectors** **very** **important** **as** **our** **traditional** **source** **of** **performing** **arts** **in** **health**, **medical** **research**, **education** **and** **the** **social** **services**," **Robert** **Franklin** **President** **Wilson-Randolph-Hoover** **said** **Monday** **in** **announcing** **the** **grants**. **Lincoln** **Center's** **President** **Robert** **Franklin** **will** **receive** **\$400,000** **each**. **The** **Metropolitan** **Opera** **Co.** **and** **New** **York** **Philharmonic** **will** **receive** **\$400,000** **each**. **The** **Juilliard** **School**, **where** **music** **and** **the** **performing** **arts** **are** **taught**, **will** **get** **\$200,000**. **The** **Wilson-Randolph-Hoover** **Foundation**, **a** **leading** **operating** **office** **of** **Lincoln** **Center** **Consolidated** **Corporate** **Fund**, **will** **make** **its** **annual** **donation** **\$100,000** **donative**.

© Eric Xing @ CMU, 2006-2009



Reading
a noun
➤ (vs verb)

[Rustandi et al.,
2005]



Paradigms of Machine Learning

- **Supervised Learning**

- Given $D = \{X_i, Y_i\}$, learn $f(\cdot): Y_i = f(X_i)$, s.t. $D^{\text{new}} = \{X_j\} \Rightarrow \{Y_j\}$

- **Unsupervised Learning**

- Given $D = \{X_i\}$, learn $f(\cdot): Y_i = f(X_i)$, s.t. $D^{\text{new}} = \{X_j\} \Rightarrow \{Y_j\}$

- **Reinforcement Learning**

- Given $D = \{\text{env, actions, rewards, simulator/trace/real game}\}$

learn policy: $e, r \rightarrow a$, s.t. $\{\text{env, new real game}\} \Rightarrow a_1, a_2, a_3 \dots$
 utility: $a, e \rightarrow r$

- **Active Learning**

- Given $D \sim G(\cdot)$, learn $D^{\text{new}} \sim G'(\cdot)$ and $f(\cdot)$, s.t. $D^{\text{all}} \Rightarrow G'(\cdot), \text{policy}, \{Y_j\}$

Theories of Learning



For the learned $F(; \theta)$

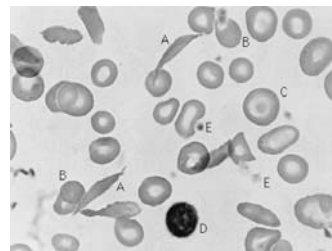
- Consistency (value, pattern, ...)
- Bias versus variance
- Sample complexity
- Learning rate
- Convergence
- Error bound
- Confidence
- Stability
- ...

© Eric Xing @ CMU, 2006-2009

Function Approximation



- **Setting:**
 - Set of possible instances X
 - Unknown target function $f: X \rightarrow Y$
 - Set of function hypotheses $H = \{h \mid h: X \rightarrow Y\}$
- **Given:**
 - Training examples $\{ \langle x_i, y_i \rangle \}$ of unknown target function f
- **Determine:**
 - Hypothesis $h \in H$ that best approximates f

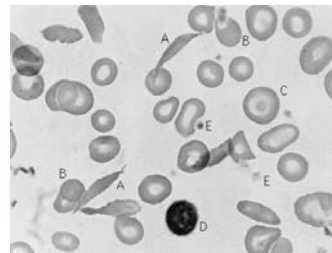
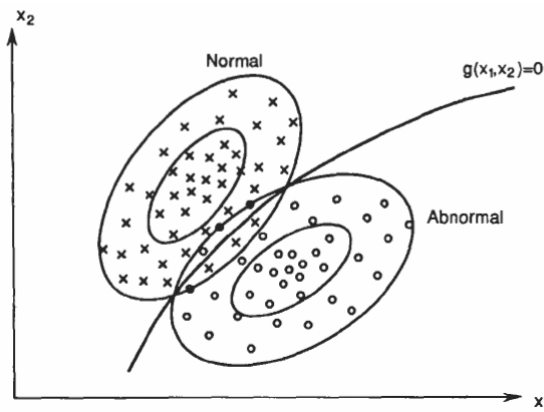


© Eric Xing @ CMU, 2006-2009

Decision-making as high-dimensional inference



- Distributions of samples from normal and abnormal machine

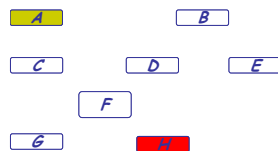


© Eric Xing @ CMU, 2006-2009

Multivariate statistical models

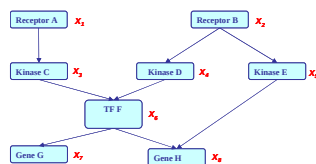


- What is a “model”?



$$P(X_1, X_2, X_3, X_4, X_5, X_6, X_7, X_8,)$$

- Graphical models:



$$\begin{aligned} &P(X_1, X_2, X_3, X_4, X_5, X_6, X_7, X_8) \\ &= P(X_1) P(X_2) P(X_3|X_1) P(X_4|X_2) P(X_5|X_2) \\ &\quad P(X_6|X_3, X_4) P(X_7|X_6) P(X_8|X_6, X_5) \end{aligned}$$

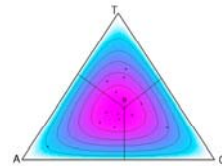
© Eric Xing @ CMU, 2006-2009



Inference Prediction Decision-Making under uncertainty

...

- Statistical Machine Learning
- Function Approximation: $F(\cdot | \theta)$?
- Density Estimation

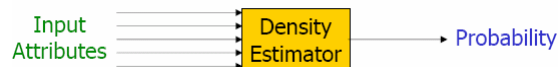


© Eric Xing @ CMU, 2006-2009



Density Estimation

- A Density Estimator learns a mapping from a set of attributes to a Probability



- Often know as *parameter estimation* if the distribution form is specified
 - Binomial, Gaussian ...
- Three important issues:
 - Nature of the data (iid, correlated, ...)
 - Objective function (MLE, MAP, ...)
 - Algorithm (simple algebra, gradient methods, EM, ...)
 - Evaluation scheme (likelihood on test data, predictability, consistency, ...)

© Eric Xing @ CMU, 2006-2009

Discrete Distributions

- Bernoulli distribution: $\text{Ber}(p)$

$$p(x) = \begin{cases} 1-p & \text{for } x=0 \\ p & \text{for } x=1 \end{cases} \Rightarrow p(x) = p^x (1-p)^{1-x}$$



- Multinomial distribution: $\text{Mult}(1, \theta)$

- Multinomial (indicator) variable:

$$X = \begin{bmatrix} X_1 \\ X_2 \\ X_3 \\ X_4 \\ X_5 \\ X_6 \end{bmatrix}, \quad \text{where} \quad \begin{aligned} &X_j \in [0,1], \text{ and } \sum_{j=1, \dots, 6} X_j = 1 \\ &X_j = 1 \text{ w.p. } \theta_j, \quad \sum_{j=1, \dots, 6} \theta_j = 1. \end{aligned}$$



$$\begin{aligned} p(x(j)) &= P\{X_j = 1, \text{ where } j \text{ index the dice-face}\} \\ &= \theta_j = \theta_A^{x_A} \times \theta_C^{x_C} \times \theta_E^{x_E} \times \theta_T^{x_T} = \prod_k \theta_k^{x_k} = \theta^x \end{aligned}$$

© Eric Xing @ CMU, 2006-2009

Continuous Distributions

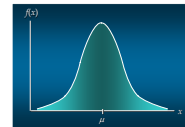
- Uniform Probability Density Function

$$p(x) = \begin{cases} 1/(b-a) & \text{for } a \leq x \leq b \\ 0 & \text{elsewhere} \end{cases}$$



- Normal (Gaussian) Probability Density Function

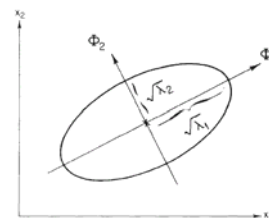
$$p(x) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(x-\mu)^2}{2\sigma^2}}$$



- The distribution is symmetric, and is often illustrated as a bell-shaped curve.
- Two parameters, μ (mean) and σ (standard deviation), determine the location and shape of the distribution.
- The highest point on the normal curve is at the mean, which is also the median and mode.
- The mean can be any numerical value: negative, zero, or positive.

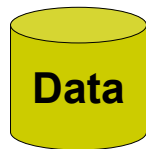
- Multivariate Gaussian

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



© Eric Xing @ CMU, 2006-2009

Density Estimation Schemes



$(x_1^{(1)}, \dots, x_n^{(1)})$
 $(x_1^{(2)}, \dots, x_n^{(2)})$
 $\dots$
 $(x_1^{(M)}, \dots, x_n^{(M)})$

Learn parameters

Algorithm

Score param

Maximum likelihood

Analytical

10^{-5}

Bayesian

Gradient

10^{-3}

Conditional likelihood

EM

10^{-15}

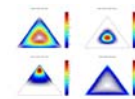
Margin

Sampling

$\dots$

$\dots$

$\dots$



© Eric Xing @ CMU, 2006-2009

Parameter Learning from *iid* Data



- Goal: estimate distribution parameters θ from a dataset of N independent, identically distributed (*iid*), fully observed, training cases

$$D = \{x_1, \dots, x_N\}$$

- Maximum likelihood estimation (MLE)

- One of the most common estimators
- With iid and full-observability assumption, write $L(\theta)$ as the likelihood of the data:

$$\begin{aligned}
 L(\theta) &= P(x_1, x_2, \dots, x_N; \theta) \\
 &= P(x_1; \theta) P(x_2; \theta), \dots, P(x_N; \theta) \\
 &= \prod_{i=1}^N P(x_i; \theta)
 \end{aligned}$$

- pick the setting of parameters most likely to have generated the data we saw:

$$\theta^* = \arg \max_{\theta} L(\theta) = \arg \max_{\theta} \log L(\theta)$$

© Eric Xing @ CMU, 2006-2009

Example: Bernoulli model



- Data:
 - We observed *N* iid coin tossing: $D=\{1, 0, 1, \dots, 0\}$
- Representation:

Binary r.v: $x_n = \{0,1\}$



- Model:
$$P(x) = \begin{cases} 1-\theta & \text{for } x=0 \\ \theta & \text{for } x=1 \end{cases} \Rightarrow P(x) = \theta^x (1-\theta)^{1-x}$$
- How to write the likelihood of a single observation x_i ?

$$P(x_i) = \theta^{x_i} (1-\theta)^{1-x_i}$$

- The likelihood of dataset $D=\{x_1, \dots, x_N\}$:

$$P(x_1, x_2, \dots, x_N | \theta) = \prod_{i=1}^N P(x_i | \theta) = \prod_{i=1}^N (\theta^{x_i} (1-\theta)^{1-x_i}) = \theta^{\sum_{i=1}^N x_i} (1-\theta)^{\sum_{i=1}^N (1-x_i)} = \theta^{\text{\#head}} (1-\theta)^{\text{\#tails}}$$

© Eric Xing @ CMU, 2006-2009

Maximum Likelihood Estimation



- Objective function:

$$\ell(\theta; D) = \log P(D | \theta) = \log \theta^{n_h} (1-\theta)^{n_t} = n_h \log \theta + (N - n_h) \log(1-\theta)$$
- We need to maximize this w.r.t. θ
- Take derivatives wrt θ

$$\frac{\partial \ell}{\partial \theta} = \frac{n_h}{\theta} - \frac{N - n_h}{1-\theta} = 0 \quad \Rightarrow \quad \hat{\theta}_{MLE} = \frac{n_h}{N} \quad \text{or} \quad \hat{\theta}_{MLE} = \frac{1}{N} \sum_i x_i$$

Frequency as sample mean

- Sufficient statistics
 - The counts, n_h , where $n_k = \sum_i x_i$, are **sufficient statistics** of data D

© Eric Xing @ CMU, 2006-2009

Overfitting

- Recall that for Bernoulli Distribution, we have

$$\hat{\theta}_{ML}^{head} = \frac{n^{head}}{n^{head} + n^{tail}}$$

- What if we tossed too few times so that we saw zero head?
We have $\hat{\theta}_{ML}^{head} = 0$, and we will predict that the probability of seeing a head next is zero!!!

- The rescue: "**smoothing**"

- Where n' is known as the pseudo- (imaginary) count

$$\hat{\theta}_{ML}^{head} = \frac{n^{head} + n'}{n^{head} + n^{tail} + n'}$$

- But can we make this more formal?

© Eric Xing @ CMU, 2006-2009

Bayesian Parameter Estimation

- Treat the distribution parameters θ also as a *random variable*
- The *a posteriori* distribution of θ after seeing the data is:

$$p(\theta | D) = \frac{p(D | \theta)p(\theta)}{p(D)} = \frac{p(D | \theta)p(\theta)}{\int p(D | \theta)p(\theta)d\theta}$$

This is Bayes Rule

$$\text{posterior} = \frac{\text{likelihood} \times \text{prior}}{\text{marginal likelihood}}$$

Bayes, Thomas (1763) An essay towards solving a problem in the doctrine of chances. *Philosophical Transactions of the Royal Society of London*, 53:370-418



The prior $p(\cdot)$ encodes our prior knowledge about the domain

© Eric Xing @ CMU, 2006-2009

Frequentist Parameter Estimation



Two people with different priors $p(\theta)$ will end up with different estimates $p(\theta|D)$.

- Frequentists dislike this “subjectivity”.
- Frequentists think of the parameter as a **fixed, unknown constant**, not a random variable.
- Hence they have to come up with different “objective” **estimators** (ways of computing from data), instead of using Bayes’ rule.
 - These estimators have different properties, such as being “unbiased”, “minimum variance”, etc.
 - The **maximum likelihood estimator**, is one such estimator.

© Eric Xing @ CMU, 2006-2009

Discussion



θ or $p(\theta)$, this is the problem!

© Eric Xing @ CMU, 2006-2009

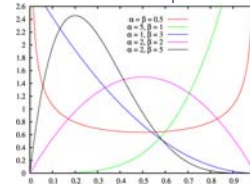
Bayesian estimation for Bernoulli



- Beta distribution:

$$P(\theta; \alpha, \beta) = \frac{\Gamma(\alpha + \beta)}{\Gamma(\alpha)\Gamma(\beta)} \theta^{\alpha-1} (1-\theta)^{\beta-1} = B(\alpha, \beta) \theta^{\alpha-1} (1-\theta)^{\beta-1}$$

- When x is discrete $\Gamma(x+1) = x\Gamma(x) = x!$



- Posterior distribution of θ :

$$P(\theta | x_1, \dots, x_N) = \frac{p(x_1, \dots, x_N | \theta) p(\theta)}{p(x_1, \dots, x_N)} \propto \theta^{n_h} (1-\theta)^{n_t} \times \theta^{\alpha-1} (1-\theta)^{\beta-1} = \theta^{n_h+\alpha-1} (1-\theta)^{n_t+\beta-1}$$

- Notice the isomorphism of the posterior to the prior,
- such a prior is called a **conjugate prior**
- α and β are hyperparameters (parameters of the prior) and correspond to the number of "virtual" heads/tails (pseudo counts)

© Eric Xing @ CMU, 2006-2009

Bayesian estimation for Bernoulli, con'd



- Posterior distribution of θ :

$$P(\theta | x_1, \dots, x_N) = \frac{p(x_1, \dots, x_N | \theta) p(\theta)}{p(x_1, \dots, x_N)} \propto \theta^{n_h} (1-\theta)^{n_t} \times \theta^{\alpha-1} (1-\theta)^{\beta-1} = \theta^{n_h+\alpha-1} (1-\theta)^{n_t+\beta-1}$$

- Maximum *a posteriori* (MAP) estimation:

$$\theta_{MAP} = \arg \max_{\theta} \log P(\theta | x_1, \dots, x_N)$$

- Posterior mean estimation:

$$\theta_{Bayes} = \int \theta p(\theta | D) d\theta = C \int \theta \times \theta^{n_h+\alpha-1} (1-\theta)^{n_t+\beta-1} d\theta = \frac{n_h + \alpha}{N + \alpha + \beta}$$

Bata parameters
can be understood
as pseudo-counts

- Prior strength: $A = \alpha + \beta$

- A can be interperated as the size of an imaginary data set from which we obtain the **pseudo-counts**

© Eric Xing @ CMU, 2006-2009

Effect of Prior Strength



- Suppose we have a uniform prior ($\alpha=\beta=1/2$), and we observe $\vec{n} = (n_h = 2, n_t = 8)$

- Weak prior $A = 2$. Posterior prediction:

$$p(x = h | n_h = 2, n_t = 8, \vec{\alpha} = \vec{\alpha} \times 2) = \frac{1+2}{2+10} = 0.25$$

- Strong prior $A = 20$. Posterior prediction:

$$p(x = h | n_h = 2, n_t = 8, \vec{\alpha} = \vec{\alpha} \times 20) = \frac{10+2}{20+10} = 0.40$$

- However, if we have enough data, it washes away the prior. e.g., $\vec{n} = (n_h = 200, n_t = 800)$. Then the estimates under weak and strong prior are $\frac{1+200}{2+1000}$ and $\frac{10+200}{20+1000}$, respectively, both of which are close to 0.2

© Eric Xing @ CMU, 2006-2009

Example 2: Gaussian density



- Data:
 - We observed N iid real samples:
 $D = \{-0.1, 10, 1, -5.2, \dots, 3\}$

- Model: $P(x) = (2\pi\sigma^2)^{-1/2} \exp\{-(x - \mu)^2 / 2\sigma^2\}$

- Log likelihood:

$$\ell(\theta; D) = \log P(D | \theta) = -\frac{N}{2} \log(2\pi\sigma^2) - \frac{1}{2} \sum_{n=1}^N \frac{(x_n - \mu)^2}{\sigma^2}$$

- MLE: take derivative and set to zero:

$$\begin{aligned} \frac{\partial \ell}{\partial \mu} &= (1/\sigma^2) \sum_n (x_n - \mu) & \Rightarrow & \mu_{MLE} = \frac{1}{N} \sum_n (x_n) \\ \frac{\partial \ell}{\partial \sigma^2} &= -\frac{N}{2\sigma^2} + \frac{1}{2\sigma^4} \sum_n (x_n - \mu)^2 & & \sigma_{MLE}^2 = \frac{1}{N} \sum_n (x_n - \mu_{MLE})^2 \end{aligned}$$

© Eric Xing @ CMU, 2006-2009

MLE for a multivariate-Gaussian



- It can be shown that the MLE for μ and Σ is

$$\mu_{MLE} = \frac{1}{N} \sum_n (x_n)$$

$$\Sigma_{MLE} = \frac{1}{N} \sum_n (x_n - \mu_{ML})(x_n - \mu_{ML})^T = \frac{1}{N} S$$

where the scatter matrix is

$$S = \sum_n (x_n - \mu_{ML})(x_n - \mu_{ML})^T = \left(\sum_n x_n x_n^T \right) - N \mu_{ML} \mu_{ML}^T$$

- The sufficient statistics are $\sum_n x_n$ and $\sum_n x_n x_n^T$.
- Note that $X^T X = \sum_n x_n x_n^T$ may not be full rank (eg. if $N < D$), in which case Σ_{ML} is not invertible

$$x_n = \begin{pmatrix} x_n^1 \\ x_n^2 \\ \vdots \\ x_n^K \end{pmatrix}$$

$$X = \begin{pmatrix} x_1^T \\ x_2^T \\ \vdots \\ x_N^T \end{pmatrix}$$

© Eric Xing @ CMU, 2006-2009

Bayesian estimation



- Normal Prior:

$$P(\mu) = (2\pi\sigma_0^2)^{-1/2} \exp\left\{-\frac{(\mu - \mu_0)^2}{2\sigma_0^2}\right\}$$

- Joint probability:

$$P(x, \mu) = (2\pi\sigma^2)^{-N/2} \exp\left\{-\frac{1}{2\sigma^2} \sum_{n=1}^N (x_n - \mu)^2\right\} \\ \times (2\pi\sigma_0^2)^{-1/2} \exp\left\{-\frac{(\mu - \mu_0)^2}{2\sigma_0^2}\right\}$$

- Posterior:

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

$$\text{where } \tilde{\mu} = \frac{N/\sigma^2}{N/\sigma^2 + 1/\sigma_0^2} \bar{x} + \frac{1/\sigma_0^2}{N/\sigma^2 + 1/\sigma_0^2} \mu_0, \text{ and } \tilde{\sigma}^2 = \left(\frac{N}{\sigma^2} + \frac{1}{\sigma_0^2} \right)^{-1}$$

© Eric Xing @ CMU, 2006-2009

Bayesian estimation: unknown μ , known σ



$$\mu_N = \frac{N/\sigma^2}{N/\sigma^2 + 1/\sigma_0^2} \bar{x} + \frac{1/\sigma_0^2}{N/\sigma^2 + 1/\sigma_0^2} \mu_0, \quad \tilde{\sigma}^2 = \left(\frac{N}{\sigma^2} + \frac{1}{\sigma_0^2} \right)^{-1}$$

- The posterior mean is a convex combination of the prior and the MLE, with weights proportional to the relative noise levels.
- The precision of the posterior $1/\sigma_N^2$ is the precision of the prior $1/\sigma_0^2$ plus one contribution of data precision $1/\sigma^2$ for each observed data point.

- Sequentially updating the mean

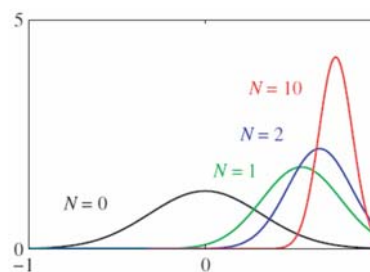
- $\mu^* = 0.8$ (unknown), $(\sigma^2)^* = 0.1$ (known)

- Effect of single data point

$$\mu_1 = \mu_0 + (x - \mu_0) \frac{\sigma_0^2}{\sigma^2 + \sigma_0^2} = x - (x - \mu_0) \frac{\sigma_0^2}{\sigma^2 + \sigma_0^2}$$

- Uninformative (vague/ flat) prior, $\sigma_0^2 \rightarrow \infty$

$$\mu_N \rightarrow \mu_0$$



© Eric Xing @ CMU, 2006-2009

Summary



- Machine Learning is Cool and Useful!!

- Functional approx.
- Multivariate statistical models
- High-dimensional inference

- Learning scenarios:

- Data
- Objective function
- Frequentist and Bayesian
- Supervised versus unsupervised,

- Density estimation

- Typical discrete distribution
- Typical continuous distribution
- Conjugate priors

© Eric Xing @ CMU, 2006-2009