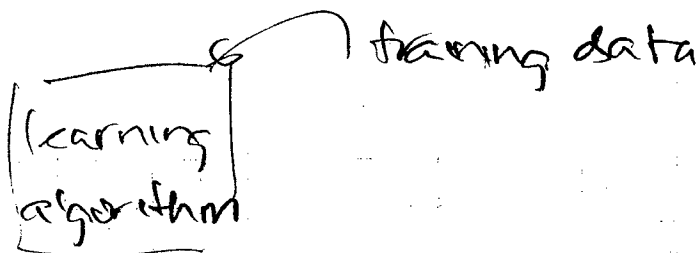


1.15.2009

Kinds of problems

- estimation
- classification (figuring out a class from which data comes)
 ↑
 e.g. classifying written digits
- regression (predicting a real number for data points)
 (e.g. predicting probability that a person gets in an accident based on age, income, # previous accidents)

classification and regression are examples of supervised learning.



1. Train the algorithm using training data
2. Use the algorithm on real data.
 trained

raining	temp ≤ 32	wind > 10 mph	?	class windows
yes	no	no	no	yes
no	yes	no	yes	yes
yes	no	yes	no	no
		⋮		

It's supervised because training examples contain the answers (output classes).

We want the generalization (true) error to be small.

↑ error on the real data

Along the way, we will be minimizing the training error.

We want the training data and the real data to be related. We are going to assume that there is some distribution on the data and the training data and the real data are independent samples from the distribution.

Classification

$(x_1, y_1), (x_2, y_2), (x_3, y_3), \dots, (x_n, y_n)$

The training data are independent (from some distribution, D)

After learning A using training data, what is the, e.g., expected error for (x, y) from D ?

The only difference between classification and regression is that y is discrete for classification and continuous for regression.

Prof. will not be providing rigorous proofs of error measures

Unsupervised Learning

Training data is $x_1, x_2, x_3, \dots, x_n$

Goal is to cluster data into groups.

Basic Probability Theory (finite probability spaces first)

Ω is called the sample space.

For now, $|\Omega| < \infty$

The elements of Ω are called atomic events.
The subsets of Ω are called events.

$$P: \mathcal{Z}^\Omega \rightarrow \mathbb{R}_{\geq 0}$$

all subsets
of Ω

$$P(\Omega) = 1$$

If $A_1, \dots, A_k$ are disjoint then $P(A_1 \cup \dots \cup A_k) = \sum P(A_i)$

Of course, it is enough to know the probs. of the atomic events (e.g. $P(\{x\})$ for all $x \in \Omega$)

$$\text{Then } P(A) = \sum_{x \in A} P(\{x\})$$

Throw a die twice

$$\Omega = \{1, \dots, 6\} \times \{1, \dots, 6\}$$

$$P(x) = \frac{1}{36}, \text{ e.g. } P((1,1)) = \frac{1}{36}$$

Two events $A, B \subseteq \Omega$ are independent if

$$P(A \cap B) = P(A) P(B)$$

$$A = \{(1, x), x \in \{1, \dots, 6\}\}$$

$$B = \{(x, 2), x \in \{1, \dots, 6\}\}$$

$$A \cap B = \{(1, 2)\}$$

$$P(A) = P(B) = \frac{1}{6}, P(A \cap B) = \frac{1}{36}$$

$$A = \{(x, y) : x + y = 7\}$$

$$B = \{(x, 2) : x \in \{1, \dots, 6\}\}$$

$$P(A) = \frac{1}{6}$$

$$A \cap B = \{(5, 2)\}$$

$$P(B) = \frac{1}{6}$$

$$P(A \cap B) = \frac{1}{36}$$

In the above case, A and B are independent
They wouldn't be if $A = \{\dots, x + y = 3\}$

In the later case, $P(A) = \frac{1}{18}$ and $P(A)P(B) \neq P(A \cap B)$

Conditional Probability

$$P(B|A) = \frac{P(A \cap B)}{P(A)} = \frac{\frac{1}{36}}{\frac{1}{18}} = \frac{1}{2}$$

$$\text{Independent} \Leftrightarrow P(B|A) = P(B)$$

Random Variables

A random variable is a function

$$X: \Omega \rightarrow \mathbb{R}$$

Throw two dice. If sum is 7, you pay \$7,
otherwise I pay \$1.
(e.g. $X((1,6)) = 7$, $X((1,2)) = -1$)

The expected value of a variable X is

$$E[X] = \sum_{w \in \Omega} P(w) X(w)$$

$$\text{e.g. } 7 \cdot \frac{1}{6} - 1 \cdot \frac{5}{6} = \frac{1}{3}$$

Theorem (Weak Law of Large Numbers)

For any $\epsilon > 0$ and $X_1, \dots, X_n$ independent, identically distributed (i.i.d),

$$\lim_{n \rightarrow \infty} P\left(\left| \frac{X_1 + \dots + X_n}{n} - E[X] \right| > \epsilon\right) = 0$$

Entropy of a (discrete) random variable X is

$$\text{entropy}(X) = - \sum_a P(X=a) \cdot \log_2 P(X=a)$$

Random variable $\rightarrow$ Event

X

$$\{w \mid X(w) = a\}$$

$$\text{e.g. } \begin{cases} X=7, \{(1,6), \dots\} \\ X=-1, \Omega - \{(1,6), \dots\} \end{cases}$$

6
 If $P(X=1) = 1/2$, $P(X=0) = 1/2$, then
 entropy $(X) = 1$

For $P(Y=0) = P(Y=1) = P(Y=2) = P(Y=3) = 1/4$,
 entropy $(Y) = 2$

If entropy $(X) = a$ and entropy $(Y) = b$,
 X and Y independent

Two r.v. are independent if all events
 created from them are independent. For
 example $(X \leq c)$ and $(Y \leq d)$ are
 independent for every c, d .

$Y \backslash X$	$X=X_1$	$X=X_2$	...	$X=X_k$
$Y=Y_1$		$P(X=X_2, Y=Y_1) = P(X=X_2) P(Y=Y_1)$		
		Independent		
$Y=Y_k$				

entropy $((X, Y)) =$

$$-\sum_{i,j} P(X=X_i \wedge Y=Y_j) \log_2 P(X=X_i \wedge Y=Y_j)$$

$$p_i = P(X = x_i)$$

$$q_j = P(Y = y_j)$$

$$= - \sum_{i=1}^k \sum_{j=1}^l p_i q_j \log_2 p_i q_j$$

$$= - \sum_{i=1}^k \sum_{j=1}^l (p_i q_j \log_2 p_i + p_i q_j \log_2 q_j)$$

$$= - \sum_{i=1}^k p_i \log_2 p_i - \sum_{j=1}^l q_j \log_2 q_j = a + b$$

$$\text{So, entropy}((X, Y)) = \text{entropy}(X) + \text{entropy}(Y).$$

Probability Space

- ① set Ω (sample space)
 B a collection of subsets of Ω
- ② function $P: B \rightarrow [0, 1]$ $B = 2^\Omega$ we took all subsets when $|\Omega| < \infty$
- ③ $P(\Omega) = 1$

$A_1, \dots, A_n, \dots$ disjoint sets from Ω

$$P\left(\bigcup_{i=1}^{\infty} A_i\right) = \sum_{i=1}^{\infty} P(A_i)$$

← before we did this for finite collections

B is a nice collection of subsets

if $\Omega = \mathbb{R}$

then

- B contains all intervals
- is closed under complement

σ -algebra

σ -field

$$A \in B \Rightarrow \mathbb{R} \setminus A \in B$$

- is closed under countable unions

$$A_1, \dots, A_n, \dots \in B \Rightarrow \bigcup_{i=1}^{\infty} A_i \in B$$

a collection of sets that satisfy the 2 properties

DEF

A random variable is a function $X: \Omega \rightarrow \mathbb{R}$
¹
 measurable

DEF

The distribution function of X is

$$f: \mathbb{R} \rightarrow [0, 1] \quad f(a) = P(X \leq a)$$

Example:

$$\Omega = \{1, 2, 3\}$$

$$P(\{1\}) = \frac{1}{4}$$

$$P(\{2\}) = \frac{1}{4}$$

$$P(\{3\}) = \frac{1}{2}$$

$$X(1) = 1$$

$$X(2) = 3$$

$$X(3) = 7$$

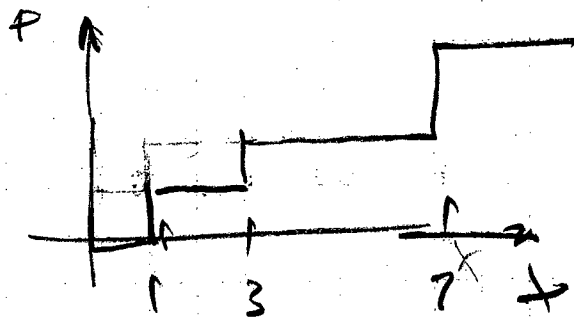
What is the distribution function of X ?

$$P(X \leq 0) = 0$$

$$P(X \leq 1) = \frac{1}{4}$$

$$P(X \leq 0.999) = 0$$

$$P(X \leq 5) = \frac{1}{2}$$

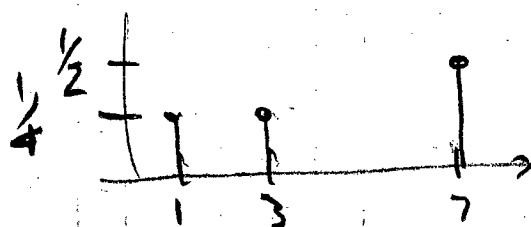


This is also called a cumulative density function if X is continuous

We will mainly encounter r.v. that are

- ① Discrete r.v. have countably many possible values

Probability mass function of x
 $f(a) = P(X=a)$



- ② Continuous r.v. (e.g. random real number from $[0, 1]$) with distr. function F

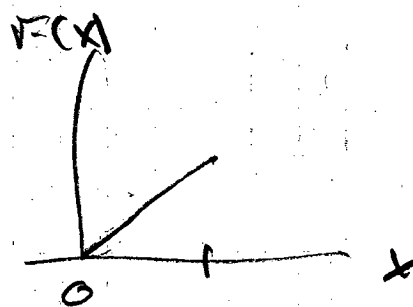
A random variable is continuous if there exists $f: \mathbb{R} \rightarrow \mathbb{R}$ such that

$$F(x) = \int_{-\infty}^x f(x) dx$$

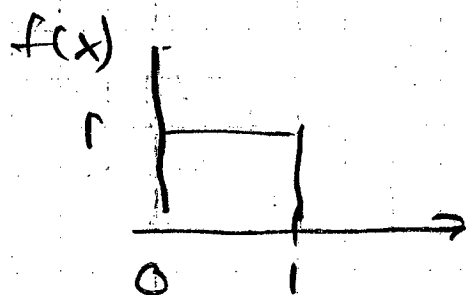
(f is density of x)

$$= P(X \leq x)$$

E.g. uniform
 $P(X \leq \frac{1}{2}) = \frac{1}{2}$
 $P(X \leq 1) = 1$



The density is

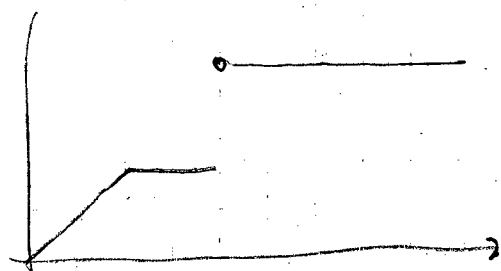


For a uniform r.v. $\in [0, 2]$, the density would be $1/2$.

This is neither discrete nor continuous

$X \sim \frac{1}{2}$ random from $[0, 1]$
 $\frac{1}{2}$ 2

because it's both



Basic Discrete Distributions

- Bernoulli Distribution (p)

$$P(X=1) = p$$

$$P(X=0) = 1-p$$

$$E[X] = p$$

- Binomial Distribution (n, p)

$$P(X=k) = \binom{n}{k} p^k (1-p)^{n-k}$$

$$k \in \{0, \dots, n\}$$

Bernoulli is a special case of binomial.

If $X_1, \dots, X_n$ are independent from Bernoulli (p), then $X = X_1 + \dots + X_n$ is from Binomial (n, p)

$$E[X+Y] = E[X] + E[Y]$$

linearity of expectation

Observe Bernoulli. How many coin flips are needed to get 1 head?

• Geometric Distribution

$$P(X=k) = p(1-p)^{k-1}$$

Let $X_1, \dots, X_n, \dots$ be independent r.v. from Bernoulli (p) .

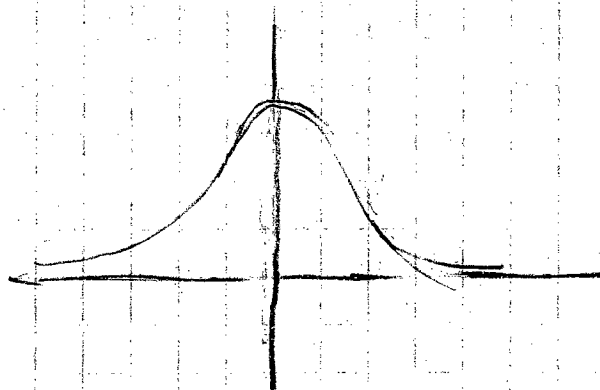
$$X = \min \{t \mid X_t = 1\}$$

$$E(X) = \frac{1}{p}$$

• Gaussian (Normal) Distribution (μ, σ)

density function =

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



Inequalities for R.V.s

• Markov's Inequality

Let X be a non-negative r.v. (i.e. $X \geq 0$)

$$\forall t, P(X \geq tE[X]) \leq \frac{1}{t} \quad (t > 0)$$

Proof:

If $P(X \geq tE[X]) > \frac{1}{t}$, then

$$E[X] = \int_{-\infty}^{\infty} x f(x) dx \geq \int_{tE[X]}^{\infty} x f(x) dx$$

$$\geq \int_{tE[X]}^{\infty} tE[X] f(x) dx = tE[X] \int_{tE[X]}^{\infty} f(x) dx > E[X]$$

$$= P(X \geq tE[X]) > \frac{1}{t}$$

So, we have $E[X] > E[X]$, which is a contradiction.

DEF

Let X be a random variable. The variance is

$$V[X] = E[(X - E[X])^2]$$

For Bernoulli:

	X	$(X - E[X])^2$
p	1	$(1-p)^2$
$1-p$	0	p^2

$$\begin{aligned} V[X] &= p(1-p)^2 + (1-p)p^2 \\ &= p(1-p) \end{aligned}$$

If $X_1, \dots, X_n$ are independent, then

$$V(X_1 + \dots + X_n) = \sum_{i=1}^n V(X_i)$$

$$V(\alpha X) = \alpha^2 V(X)$$

• Chebyshev's Inequality

$$X, E(X) < \infty, V(X) < \infty$$

$$P(|X - E(X)| \geq t) \leq \frac{V(X)}{t^2}$$

$$\text{Let } Y = (X - E(X))^2 \quad E(Y) = V(X)$$

$$P(Y \geq t E(Y)) \leq \frac{1}{t}$$

$$P((X - E(X))^2 \geq t V(X)) \leq \frac{1}{t}$$

$$\text{Let } t = \frac{1}{\epsilon^2} \frac{1}{V(X)}$$

$$P((X - E(X))^2 \geq \frac{1}{\epsilon^2}) \leq \frac{V(X)}{1/\epsilon^2}$$

Weak Law of Large Numbers

$X_1, \dots, X_n$ are independent, identically distributed ($V(X_i) < \infty, E(X_i) = \mu$)

$$\forall \epsilon > 0$$

$$\lim_{n \rightarrow \infty} P\left(\left|\frac{X_1 + \dots + X_n}{n} - \mu\right| \geq \epsilon\right) = 0$$

Proof: Apply Chebyshev to random variable which is sum here. Variance goes to 0.

$$V\left(\frac{X_1 + \dots + X_n}{n}\right) = V(X)/n$$

Chernoff - Hoeffding

$X_1, \dots, X_n$ are independent from Bernoulli. (p)

$$P\left(\left|\frac{X_1 + \dots + X_n}{n} - p\right| \geq \varepsilon\right) \leq 2e^{-2n\varepsilon^2}$$

Chebyshev would give us

$$P\left(\left|\frac{X_1 + \dots + X_n}{n} - p\right| \geq \varepsilon\right) \leq \frac{p(1-p)}{\varepsilon^2 n}$$

Chernoff - Hoeffding gives us a stronger result.

We use these things for estimation.
An example would be polling, where Bernoulli samples are votes.

Given $X_1, \dots, X_n$ from Bernoulli, the goal is to estimate the unknown p .

$$\text{Estimate} = \frac{X_1 + \dots + X_n}{n}$$

How good is this estimate?

Central Limit Theorem

Let $X_1, \dots, X_n$ be i.i.d. r.v. with mean μ and variance σ^2 .

$$\frac{\sqrt{n}}{\sigma} \left(\frac{X_1 + \dots + X_n}{n} - \mu \right) = Z_n \quad \text{let}$$

$$\lim_{n \rightarrow \infty} \sup_{a < b} \left| P(Z_n \leq a) - \int_{-\infty}^a \frac{1}{\sqrt{2\pi}} e^{-x^2/2} dx \right| = 0$$

1.22.2009

Parameter Estimation $X_1, \dots, X_n$ i.i.d from $D(\alpha)$; α is unknownThe goal is to determine α .
distribution ϵ precision
 δ confidence $f(X_1, \dots, X_n) \rightarrow G$ r.vwe want $P(|G - \alpha| > \epsilon) \leq \delta$ Maximum Likelihood $P(X_1, \dots, X_n \mid \text{assuming the parameter is } \alpha)$ likelihoodPick α for which expression is largest.ExampleAssume that the family of distributions considered is Bernoulli (α) $\alpha \in [0, 1]$ $n=6$ 0, 0, 1, 1, 0, 1

$$P(X_1=0, X_2=0, \dots, X_5=1 \mid \alpha) = (1-\alpha)^3 \alpha^3$$

The expression above is maximized for $\alpha = 1/2$.Assume the family $N(\mu, \sigma) = \frac{1}{\sigma\sqrt{2\pi}} e^{-(x-\mu)^2/2\sigma^2}$ Observed $X_1, \dots, X_n$ we want to pick μ, σ that maximize the likelihood.

The likelihood is

$$f(\mu, \sigma) = \prod_{i=1}^n \frac{1}{\sigma \sqrt{2\pi}} e^{-\frac{(x_i - \mu)^2}{2\sigma^2}}$$

Maximizing log is same as maximizing f .

$$\log f(\mu, \sigma) = \sum_{i=1}^n -\frac{(x_i - \mu)^2}{2\sigma^2} - n \ln \sigma - n \ln \sqrt{2\pi}$$

$$\frac{\partial}{\partial \mu} \log f(\mu, \sigma) = \sum_{i=1}^n \frac{(x_i - \mu)}{\sigma^2} = 0$$

$$\mu_{ML} = \frac{1}{n} \sum_{i=1}^n x_i$$

$$\frac{\partial}{\partial \sigma} \log f(\mu, \sigma) = \sum_{i=1}^n \frac{(x_i - \mu)^2}{\sigma^3} - \frac{n}{\sigma}$$

$$\sigma_{ML}^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \mu_{ML})^2$$

Assume that data actually come from $N(\mu, \sigma)$

What will μ_{ML}, σ_{ML} give us?

$$E[\mu_{ML} | x_1, \dots, x_n \sim N(\mu, \sigma)] = \mu$$

$$E[\sigma_{ML}^2 | x_1, \dots, x_n \sim N(\mu, \sigma)] =$$

without loss of generality, we can assume $\mu = 0$

$$E[\sigma_{ML}^2 | \dots] = \frac{1}{n} E\left[\sum_{i=1}^n \left(x_i - \frac{\sum x_i}{n}\right)^2 | \dots\right]$$

$$\frac{1}{n} E \left[\sum_i x_i^2 - 2 \sum_i x_i \left(\frac{\sum_j x_j}{n} \right) + \frac{1}{n^2} \left(\sum_j x_j \right)^2 \mid \dots \right] =$$

$$V(X) = E[X^2] - E[X]^2 \quad E[X_i X_j] = \rho \text{ for } i \neq j$$

$$\frac{1}{n} (n\sigma^2 - 2\sigma^2 + \sigma^2) = \frac{n-1}{n} \sigma^2$$

Since $\sigma_{ML} \neq \sigma$, we say that σ_{ML} is a biased estimator.

But it is asymptotically unbiased since

$$\lim_{n \rightarrow \infty} \sigma_{ML} = \sigma$$

Bayes Theorem

$$P(A|B) = \frac{P(A \cap B)}{P(B)}$$

$$P(B|A) = \frac{P(A \cap B)}{P(A)}$$

$$P(A|B) = \frac{P(B|A) P(A)}{P(B)}$$

We will use Bayes Theorem to calculate the parameters of a distribution.

A = distribution parameters

B = data

A prior is a belief on how likely p is.
 We observe some data.
 We update the belief based on the data.
 The updated belief is called the posterior.

$$P(A|B) = \frac{P(B|A) P(A)}{P(B)}$$

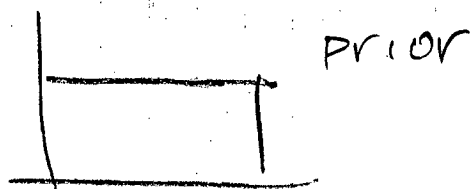
prior
normalizing factor

likelihood

posterior

Example

Assume Bernoulli distribution.



observe ϕ

posterior?

$$P(\text{parameter is } x \mid \text{observed } \phi) = \frac{(1-x)(1)}{?} = \frac{1-x}{1/2}$$

pick so that we get a probability distr.

$$P(B) = \int P(B|A) P(A) dA$$

$$P(\text{observing } 0 \mid \text{parameter is } x) = 1-x$$

$$P(\text{observing } 0) = \int (1-x) dx$$

Typically, we pick priors for mathematical convenience. We want to insure that posterior distribution is from same family as prior.

DEF

Let F be a family of likelihoods.

A family of distributions G is conjugate prior to F when

$$\text{prior} \in F, \text{likelihood} \in G \Rightarrow \text{posterior} \in F$$

What would be a nice family of priors for Bernoulli?

$$P(X_1 = a_1, \dots, X_n = a_n \mid \text{parameter is } p)$$

$$a_1, \dots, a_n \in \{0, 1\}$$

$$P(\dots) = (1-p)^{n-\sum a_i} p^{\sum a_i}$$

Prior from a distribution

$$P(z) \propto (1-z)^A z^B$$

Posterior is from

$$\begin{aligned} P(z \mid X_1 = a_1, \dots, X_n = a_n) &\propto (1-z)^A z^B (1-z)^{n-\sum a_i} z^{\sum a_i} \\ &= (1-z)^{A'} z^{B'} \end{aligned}$$

$$A' = A + n - \sum a_i$$

$$B' = B + \sum a_i$$

Beta (A, B) distribution is a distribution on $[0, 1]$

$$P(z) \propto z^{A-1} (1-z)^{B-1} \quad \text{continuous}$$

The Beta distribution is the conjugate prior for the Bernoulli distribution.
discrete

Let's look at Gaussian distributions. Let's first look at μ , assuming σ is fixed.

$$P(x_1, \dots, x_n | \mu) =$$

We are looking for priors over μ . The likelihood function is

$$P(x_1, \dots, x_n | \mu) = \prod_{i=1}^n \frac{1}{\sigma \sqrt{2\pi}} e^{-\frac{(x_i - \mu)^2}{2\sigma^2}}$$

We need something that has a quadratic function as an exponent of e .

The prior on μ is $N(\mu_0, \sigma_0)$

Bayesian "setting"

prior (distribution
on the
parameters)

likelihood (from data)

posterior (different distribution
on the data)

parameter

e.g. $x_1, \dots, x_n$ are from $\text{Bernoulli}(p)$

We want prior $\times$ likelihood = posterior
such that posterior is from same distribution
as prior.

prior $p \sim \text{Beta}(a, b)$

$$P(p=x) \propto x^{a-1} (1-x)^{b-1}$$

observed $d_1, \dots, d_n \in \{0, 1\}$

$$P(0 | p=x) = 1-x$$

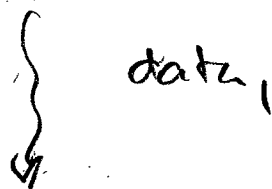
$$P(1 | p=x) = x$$

$$P(d_i | p=x) = x^{d_i} (1-x)^{1-d_i}$$

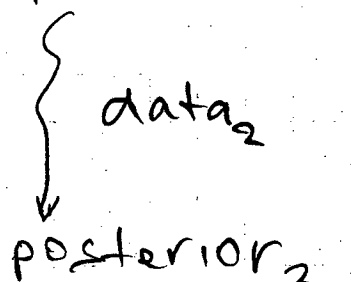
$$\text{likelihood} = x^{\sum d_i} (1-x)^{n - \sum d_i}$$

$$\text{posterior} \sim \text{Beta}(a + \sum d_i, b + n - \sum d_i)$$

prior



posterior₁ = use as prior for next batch of data



Is this the same

prior $\longrightarrow$ posterior
data₁ + data₂

Is this the same as getting all of the data at once? Yes

This is called sequential approach to Bayesian.

Beta(a, b)

$\downarrow d_1, \dots, d_n$

Beta(a + $\sum d_i$, b + n - $\sum d_i$)

$\downarrow d'_1, \dots, d'_m$

Beta(...)

Example 2

Our model of the world is Gaussian

$$x_1, \dots, x_n \sim N(\mu, \sigma^2)$$

We want to estimate $\lambda = \frac{1}{\sigma^2}$

In this case the conjugate prior is the gamma distribution

$$\text{prior} \sim \text{Gamma}(a, b) \propto \lambda^{a-1} e^{-b\lambda}$$

$$\text{prior} \sim \text{Gamma}(a, b)$$

$$\downarrow$$

data
 $x_1, \dots, x_n$

$$\text{posterior} \sim \text{Gamma}(a', b')$$

$$a' = a + \frac{n}{2}$$

$$b' = b + \frac{1}{2} \sum_{i=1}^n (x_i - \mu)^2 \quad \mu = \frac{1}{n} \sum_{i=1}^n x_i$$

In this case, processing data all at once yields different results than sequential.
consider adding one point at a time.
b would not change in sequential case.

Parametric Framework

We always assume data is independent

MLE:

We have some parametric model ($\theta =$ parameter)

$$\text{likelihood}(\text{data} | \theta) = P(x_1 = d_1, \dots, x_n = d_n | \theta) = \prod P(x_i = d_i | \theta)$$

Maximum likelihood method says to pick θ which maximizes likelihood.

The KL divergence of f, g is

$$D(f, g) = \int f \log \frac{f}{g} dx$$

- $D(f, g) \geq 0$
- $D(f, f) = 0$

Since $\log$ is monotonic, maximizing likelihood is same as maximizing $\log$ likelihood

Maximizing $\log$ likelihood

$$\frac{1}{n} \sum_{i=1}^n \log P(x_i = d_i | \theta)$$

What is this converging to?

$$E[\log P(x_1 = d_1 | \theta)] =$$

$$\int d_1 \sim \theta^* \quad ("d_1 \text{ comes from some distribution } \theta^*")$$

not in lecture

$$= \int P(x_1 = d_1 | \theta^*) \log P(x_1 = d_1 | \theta) dx_1$$

$$= -KL(f_{\theta^*} | f_{\theta}) + \int P(x_1 = d_1 | \theta^*) \log P(x_1 = d_1 | \theta^*) dx_1$$

We are trying to minimize the KL-divergence.

REGRESSION

$$(x_1, y_1), \dots, (x_n, y_n)$$

$$x_i \in \mathbb{R}^d$$

$$y_i \in \mathbb{R}$$

The goal is to fit a linear function to the data to predict y_i from x_i .

LSE (Least Squares Estimator)

$$\underline{x}_i = (x_{i1}, \dots, x_{id})$$

$$w_1 x_{i1} + \dots + w_d x_{id} = \underline{w}^T \underline{x}_i$$

For LSE we want to pick $\underline{w}$ to minimize

$$\sum_{i=1}^n (\underline{w}^T \underline{x}_i - y_i)^2$$

$$\begin{matrix} n \\ \left\{ \begin{array}{c} \boxed{\begin{matrix} x_1 \\ \vdots \\ x_n \end{matrix}} \end{array} \right. = B \end{matrix} \quad Y = \begin{matrix} \boxed{\begin{matrix} y_1 \\ \vdots \\ y_n \end{matrix}} \end{matrix}$$

$\underbrace{\hspace{10em}}_d$

BW

$$w^T x_1$$

27

$$Bw =$$

$$w^T x_n$$

Goal is to minimize $\|Bw - y\|_2^2$

Given $B \in \mathbb{R}^{n \times d}$

$$y \in \mathbb{R}^n$$

Find w minimizing $\|Bw - y\|_2^2$

Note that $\|u\|_2^2 = u^T u$

Minimize:

$$y^T y + w^T B^T B w - 2y^T B w$$

$$\frac{\partial}{\partial w} a^T w = a^T$$

$$a_1 w_1 + a_2 w_2 + \dots + a_d w_d$$

For $w = (w_1, \dots, w_d)$

$$\frac{\partial}{\partial w} f = \left(\frac{\partial}{\partial w_1} f, \dots, \frac{\partial}{\partial w_d} f \right)$$

$$\frac{\partial}{\partial w} w^T G w = \frac{\partial}{\partial w} \left(\sum_{i=1}^d \sum_{j=1}^d w_i w_j G_{ij} \right) = w^T (G + G^T)$$

$$G \in \mathbb{R}^{d \times d}$$

$$\frac{\partial}{\partial w_k} \left(\sum_{\substack{j=1, \\ j \neq k}}^d w_j G_{kj} + \sum_{\substack{i=1, \\ i \neq k}}^d w_i G_{ik} + 2w_k G_{kk} \right)$$

symmetric 28

So, to minimize $W^T \underbrace{B^T B}_{\text{symmetric}} W - 2y^T B W = (\dots)$

$$\frac{d}{dW} (\dots) = 2W^T B^T B - 2y^T B = 0$$

$$W^T (B^T B) = y^T B$$

If $B^T B$ is invertible,

$$W^T = y^T B (B^T B)^{-1}$$

If $B^T B$ is not invertible, then ...

$A + \epsilon I$ is singular iff $\det |A + \epsilon I| = 0$

polynomial in ϵ
of degree $\leq d$

$$L^+ B^+ = \lim_{\epsilon \rightarrow 0} B (B^T B + \epsilon I)^{-1}$$

This is called the Moore Penrose pseudo inverse.

$$W^T = y^T B^+ \quad \text{this is a "dagger", not a "plus"}$$

Let A be an $m \times n$ matrix. The SVD of A is

$$A = U D V^T$$

where $U \in \mathbb{R}^{m \times m}$, $D \in \mathbb{R}^{m \times n}$, $V \in \mathbb{R}^{n \times n}$

U, V are orthogonal, that is

$$U U^T = I, \quad V V^T = I,$$

and D is diagonal.

How to compute A^+ ?

1. $A = UDV^T$

2. Let $A^+ = U D^+ V^T$

For every $\neq 0$ entry take $1/\text{that}$

$$\begin{pmatrix} 1 & 0 & 0 \\ 0 & 2 & 0 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1/2 & 0 \end{pmatrix}$$

Assume that $B = UDV^T$, Then $B^T = VD^T U^T$

~~$$UDV^T (UDV^T V D^T U^T + \epsilon I)^{-1} =$$~~

~~$$UDV^T (U D D^T U^T + \epsilon U U^T)^{-1} =$$~~

~~Then someone said it should have been~~

$$w^T = y^T B (B^T B)^{-1}$$

$$B^+ = \lim_{\epsilon \rightarrow 0} B (B^T B + \epsilon I)^{-1}$$

$$UDV^T (VD^T U^T U D V^T + \epsilon I)^{-1} =$$

$$UDV^T V (D^T D + \epsilon I)^{-1} V^T =$$

$$UD (D^T D + \epsilon I)^{-1} V^T$$

Identities

$$(AB)^T = B^T A^T$$

$$(AB)^{-1} = B^{-1} A^{-1}$$

Motivation

Assume the model is $y = w^T x + \epsilon$

$w \in \mathbb{R}^d$, $\epsilon \in N(0, \sigma)$.

Then MLE = least squares.

Classification

regression $\rightarrow$ continuous output
classification $\rightarrow$ discrete

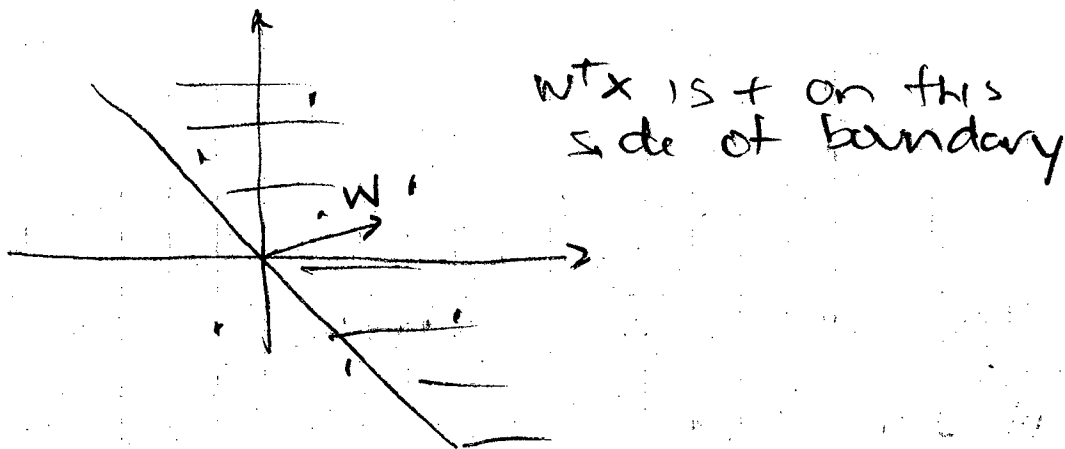
Simplest classifier is the:

Perceptron

Perceptron is a linear classifier.

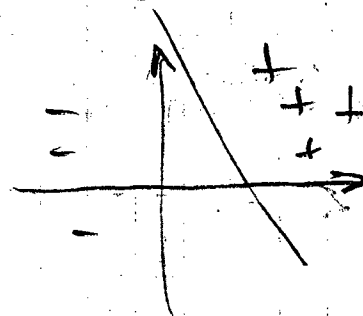
$$x = \begin{bmatrix} \end{bmatrix} \quad t \in \{+, -\} \text{ or } t \in \{1, -1\}$$

We get binary classification from $\text{sign}(w^T x)$

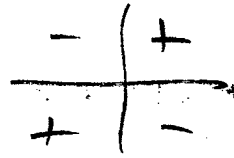


For a boundary not passing through origin
append a 1 to x , add new w_c

$$x = \begin{bmatrix} \end{bmatrix} \quad z = \begin{bmatrix} \end{bmatrix}$$

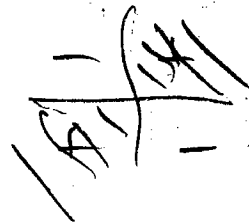


Perceptron doesn't work for

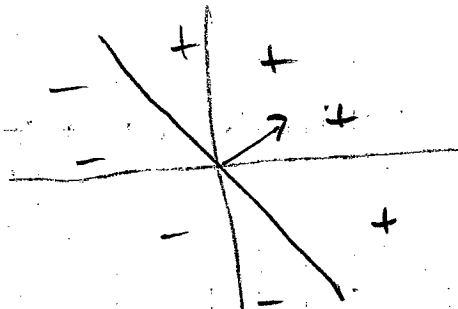


Perceptron only works if data is linearly separable.

We need to transform data



How do we pick w ?



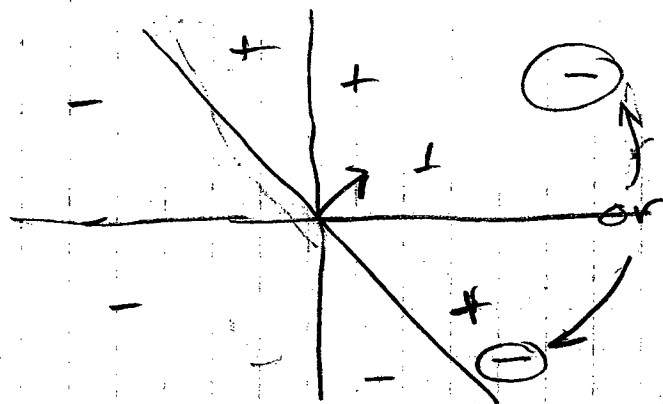
while some points misclassified
for i
if $\text{sgn}(w^T x_i) \neq t_i$
 $w \leftarrow w + \eta t_i x_i$

(η is the learning rate)

The above is called the perceptron update rule.
It will converge if data is linearly separable.

What if your data is not linearly separable?

How much do we want to penalize a weight vector for a difficult training example?



We will have an error function that tells us how bad a point is.

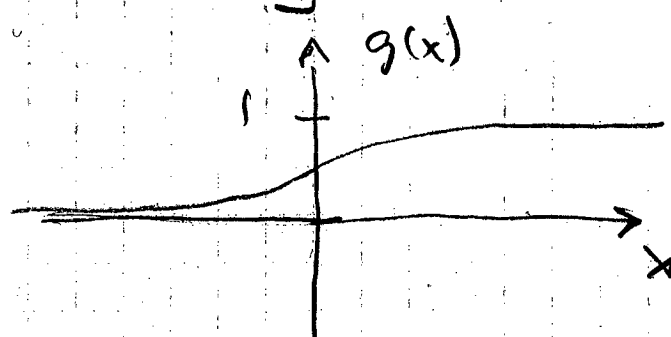
$$E(y, t) = \frac{1}{2}(t - y)^2 \quad \text{for } y = w^T x$$

⤴ This is not a good fit for classification because we don't want to respond too much for really bad data points.

Define an activation function g

$$y = g(w^T x)$$

Try $g(x) = \frac{e^x}{e^{-x} + e^x}$



Let's use $t \in \{0, 1\}$

We will use gradient descent to minimize the error

while not satisfied
for $i = 1$ to N
$$w \leftarrow w + \eta \frac{\partial E(y_i, t_i)}{\partial w}$$

$$\frac{\partial E}{\partial w} = \frac{\partial E}{\partial y} \frac{\partial y}{\partial w}$$

$$= \frac{\partial E}{\partial y} g'(w^T x) x$$

This ends up being a little more sophisticated than previous case.

"while not satisfied" $\Rightarrow$ while $|w - w_{old}| > \epsilon$

Maximum Entropy = Logistic Regression

this is good for problems with ≥ 2 labels.

You are learning the distribution

$$p(y|x)$$

and you want it to have a high entropy. You want to be "uncertain about what you can't see from the data".

$$f_i(x, y)$$

$$f = \begin{bmatrix} \end{bmatrix}$$

If x and y are both discrete, for example,

$$f_i = \begin{cases} 1 & \text{if } x = \text{tion}, y = N \\ 0 & \text{otherwise} \end{cases}$$

$$\max_p H_p(y|x) \text{ s.t. } E_{\tilde{p}}[f_i] = E_p[f_i] \quad \forall i$$

$$\forall x, \sum_y p(y|x) = 1$$

$\tilde{p}$ = empirical distribution = sample distribution

$$\tilde{p}(x, y) = \frac{n(x, y)}{N}$$

All of our features should have the same expectation in the real distribution as in the empirical distribution.

How do we solve this?

$$L = \left[\sum_{x, y} p(x, y) \log \frac{1}{p(y|x)} + \sum_i \lambda_i (p(x, y) f_i(x, y) - \tilde{p}(x, y) f_i(x, y)) \right] \\ + \sum_x \mu_x (\sum_y p(y|x) - 1)$$

Lagrangian

Whenever you have a problem in the form of

$$\max f(x) \text{ subject to the constraint } g(x) = 0,$$

you can solve by

$$0 = f'(x) + \lambda g'(x)$$

$$0 = g(x)$$

$$H(y|x) = \sum_{x,y} p(x) p(y|x) \log \frac{1}{p(y|x)}$$

We assume $p(x,y) = \tilde{p}(x) \cdot p(y|x)$

$$\frac{\partial L}{\partial p(y|x)} = \left(\log \frac{1}{p(y|x)} - 1 \right) \tilde{p}(x) + \sum_i \lambda_i \tilde{p}(x) f_i(x,y) + \mu_x = 0$$

$$\log p(y|x) = \sum_i \lambda_i \tilde{p}(x) f_i(x,y) + c$$

$$p(y|x) = \frac{1}{z(x)} e^{\sum_i \lambda_i f_i(x,y)}$$

$z(x)$ is the normalization factor

$$z(x) = \sum_y e^{\sum_i \lambda_i f_i(x,y)}$$

$$\left(\begin{aligned} & \sum_{x,y} \tilde{p}(x) p(y|x) \left[- \sum_i \lambda_i f_i(x,y) + \log z(x) \right] + \sum_i \lambda_i \tilde{p}(x) p(y|x) \cdot f_i(x,y) \\ & \sum_{x,y} \tilde{p}(x) p(y|x) \log z(x) \\ & \max_p E_p \log z(x) \end{aligned} \right)$$

This is a simplification of the original equations.

There is no closed form for λ .

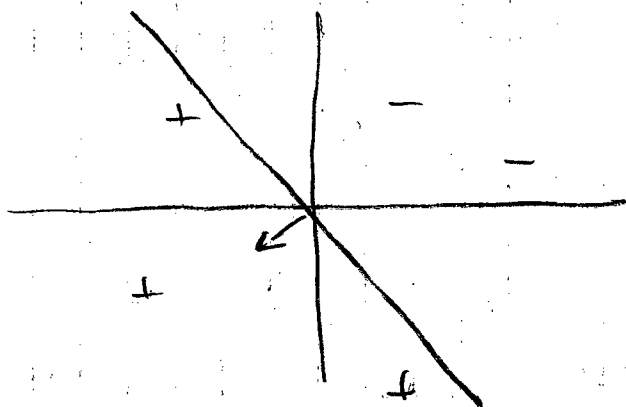
This is another way to penalize pts on wrong side of decision surface.

Tutorial W 6-7Perceptron AlgorithmInput is x_1, y_1

$$x_i \in \mathbb{R}^d$$

 x_n, y_n

$$y_i \in \{\pm 1\}$$

Goal is to learn classifier $\text{sign}(w^T x)$ 

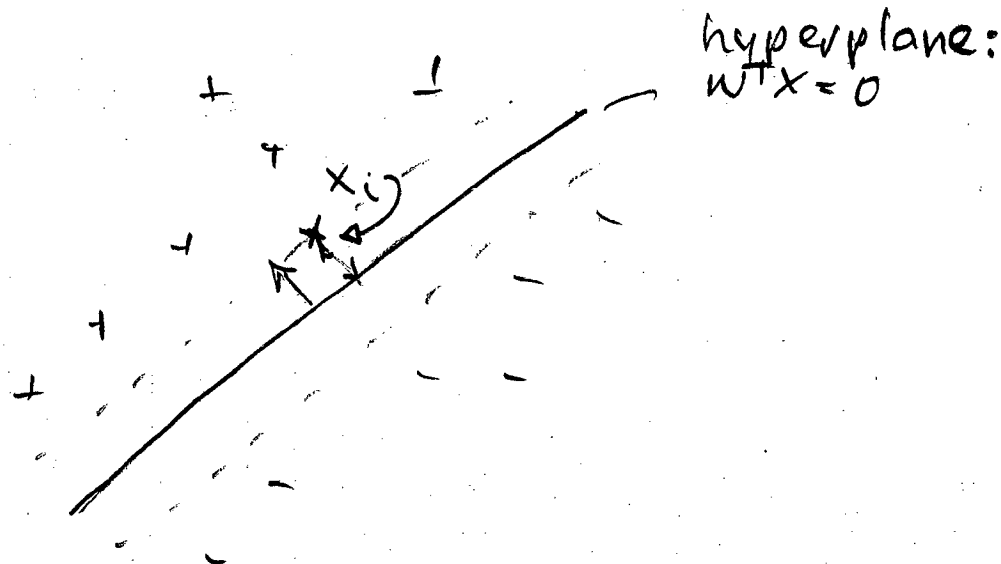
$$w \leftarrow (\phi, \dots, \phi)$$

for i from 1 to n if $\text{sign}(w^T x_i) \neq y_i$ then

$$w \leftarrow w + y_i x_i / \|x_i\|_2$$

Assume that the negative examples can be separated from the positive examples by a hyperplane (w).

Assume that no point can get "too close" to the hyperplane.



distance from hyperplane = $\frac{w^T x_i}{\|w\|_2}$

We assume the quantity above is "large" for every x_i .

Even stronger assumptions ($\exists w^*$) a "good w "

- separable
- large margin $\sigma = \min_i \frac{x_i^T w^*}{\|w^*\|_2}$ (The closest pt to hyperplane is "far" from it)

One more assumption: $\|x_i\|_2 \leq R$

The perceptron algorithm makes, at most, $O(\frac{R^2}{\sigma^2})$ updates to w , and finds the separating hyperplane w .

In the algorithm, let's assume that $\|x_i\|_2 = 1$. Then, nothing changes.

Let's also assume, w.l.o.g., that

$$\|w^*\|_2 = 1$$

all x_i are normalized to length 1.

What happens after update?

$$W' \leftarrow W + \eta_i X_i$$

$$W'^T W^* = W^T W^* + \eta_i X_i^T W^* \\ \geq W^T W^* + \frac{\sigma}{R}$$

This indicates that W' is getting closer to W^* .

This R is here because we normalized X_i .

$$\|W'\|_2^2 = W^T W + 2 \eta_i X_i^T W + \eta_i^2 X_i^T X_i$$

Because of normalization $= 1$

$$\leq 0$$

$$\therefore \|W'\|_2^2 \leq \|W\|_2^2 + 1$$

After φ updates,

$$W^T W^* \geq \varphi \cdot \frac{\sigma}{R}$$

$$W^T W^* \leq \|W\| \|W^*\| \leq \sqrt{\varphi}$$

$$\text{So, } \varphi \leq \frac{R^2}{\sigma^2}$$

$\nearrow \|W\|_2^2 \leq \varphi$ because η is by 1 each update

$\|W^*\|_2^2 = 1$ by assumption

Now, we are going to make perceptron algorithm work in a higher dimensional space.

Project $x \mapsto f(x) \in \mathbb{R}^{d'}$

Find the Separator in $\mathbb{R}^{d'}$ (instead of $\mathbb{R}^d$)

We assume that d' is big. Our goal is to run the perceptron algorithm in $\mathbb{R}^{d'}$ without computing f .

eg. $(x_1, x_2) \mapsto (x_1, x_2, x_1^2, x_1 x_2, x_2^2)$

For the perceptron algorithm

$$w = \sum_{i=1}^n \alpha_i x_i$$

on input x, y , compute

$$\text{sign}(w^T x) \stackrel{?}{=} y$$

$$\text{sign}\left(\sum_{i=1}^n \alpha_i (x_i^T x)\right) \stackrel{?}{=} y$$

We only need to look at inner products of data points to update α_i . $(x_i^T x)$

In the projected space, these become

$$f(x_i)^T f(x)$$

But we want to avoid calculating f . So, instead, we have

$$\text{DEF } K(x_i, x_j) = f(x_i)^T f(x_j)$$

$K: \mathbb{R}^d \times \mathbb{R}^d \rightarrow \mathbb{R}$ is called a kernel if

$\exists f: \mathbb{R}^d \rightarrow \mathbb{R}^\infty$ such that

$$K(x_i, x_j) = f(x_i)^T f(x_j)$$

Examples of kernels

(we drop i, j subscripts)

$$K(x, z) = x^T z \quad f \text{ is identity}$$

$$K(x, z) = (x^T z)^2$$

In 2 dimensions

$$K(x, z) = (x^T z)^2 = (x_1 z_1 + x_2 z_2)^2 \\ = x_1^2 z_1^2 + 2x_1 x_2 z_1 z_2 + x_2^2 z_2^2$$

$$\text{So, } f(x) = (x_1^2, \sqrt{2}x_1 x_2, x_2^2)$$

$$K(x, z) = e^{-\|x - z\|_2^2}$$

↑ In this case, f would go to ∞ dimensional space.

Which functions are kernels?

Mercer's Theorem states that K is a kernel iff

$$(\forall n) (\forall x_1, \dots, x_n \in \mathbb{R}^d)$$

the matrix $A = (K(x_i, x_j))$ is positive semi-definite.

A is positive semi-definite

$$A \succeq 0 \iff (\forall x) x^T A x \geq 0$$

$$(1, -1) \begin{pmatrix} 1 & 3 \\ 3 & 1 \end{pmatrix} \begin{pmatrix} 1 \\ -1 \end{pmatrix} = -4$$

So $\nearrow$ is not positive semi-definite

$$\left. \begin{array}{l} A \geq 0 \\ B \geq 0 \end{array} \right\} \Rightarrow A+B \geq 0$$

Corollary: If K_1, K_2 are kernels, then their sum $K_1 + K_2$ is a kernel.

$$\left. \begin{array}{l} A \geq 0 \\ B \geq 0 \end{array} \right\} \Rightarrow \begin{pmatrix} a_{11} & b_{11} & \dots & a_{1n} & b_{1n} \\ \vdots & \vdots & & \vdots & \vdots \\ a_{n1} & b_{n1} & \dots & a_{nn} & b_{nn} \end{pmatrix} \geq 0$$

$K(x,y) = (x^T y)^2$ is a kernel

Note that if $A \geq 0$, then $A^T A \geq 0$.

$$x^T A^T A x = \|Ax\|_2^2 \geq 0$$

If $A \geq 0$, then

$$\begin{pmatrix} e^{a_{11}} & \dots & e^{a_{1n}} \\ \vdots & & \vdots \\ e^{a_{n1}} & \dots & e^{a_{nn}} \end{pmatrix} ?$$

Note that $e^x = 1 + x + \frac{x^2}{2!} + \dots$

$$x^T \begin{pmatrix} 1 & \dots & 1 \\ \vdots & & \vdots \\ 1 & \dots & 1 \end{pmatrix} x = x^T 11^T x \geq 0$$

2.4.2009

Tutorial

Andrew Ng web page

From (2.14)

$$F[y(x) + \epsilon \eta(x)] = F[y(x)] + \epsilon \int \frac{\delta F}{\delta y(x)} \eta(x) dx + O(\epsilon^2)$$

$$\int [p(x) + \epsilon \eta(x)] \ln[p(x) + \epsilon \eta(x)] dx = \int p(x) \ln p(x) dx + \dots$$

$$\dots + \epsilon \int \left(\frac{\delta}{\delta p(x)} p(x) \ln p(x) \right) dx$$

We are not taking derivative w.r.f. $p(x)$. We are interested in the derivative above being 0 everywhere

2.5.2009

Goal: Get a good classifier which is a hyperplane (maybe in a higher dimensional space)

If the data is separable with a large margin, the perceptron algorithm "works".

For today, the goal is to find the hyperplane with the largest possible margin.

$$x^{(i)}, y^{(i)} \quad i \in \{1, \dots, n\}$$

$$y^{(i)} \in \{\pm 1\}$$

$$x^{(i)} \in \mathbb{R}^d$$

$$(w^T x^{(i)} + b) y^{(i)} \geq \gamma, \quad \max \gamma, \quad \|w\|_2^2 = 1$$

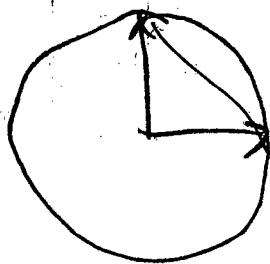
the variables are w, b, γ

We want to make this a convex optimization problem.

$$\text{Let } b' = \frac{b}{\gamma}, \quad w' = \frac{w}{\gamma}, \quad \|w'\|_2 = \frac{1}{\gamma}$$

$$\min \|w'\|_2, \quad (w'^T x^{(i)} + b') y^{(i)} \geq 1$$

Constraint is not convex, because you can draw a line between 2 w' 's that satisfy $\|w'\|_2 = 1$, but some points on the line have length < 1 .



The variables are now w', b' .

Assume w_1, b_1 and w_2, b_2 satisfy constraints.

Does $\frac{w_1 + w_2}{2}, \frac{b_1 + b_2}{2}$ satisfy the constraint?

Answer is yes.

$$\left(\frac{w_1 + w_2}{2} x^{(i)} + \frac{b_1 + b_2}{2} \right) y^{(i)} \geq 1 ?$$

$$\underbrace{\frac{1}{2} (w_1^T x^{(i)} + b_1) y^{(i)}}_{\geq 1} + \underbrace{\frac{1}{2} (w_2^T x^{(i)} + b_2) y^{(i)}}_{\geq 1} \geq 1$$

- Convex feasible set
- Objective is a convex function (minimization)

If you have these two, optimization is always possible

$$\left\| \frac{w_1 + w_2}{2} \right\|_2^2 \leq \frac{\|w_1\|_2^2 + \|w_2\|_2^2}{2}$$

$$(w_1 + w_2)^T (w_1 + w_2) \leq 2 (w_1^T w_1 + w_2^T w_2)$$

$$0 \leq w_1^T w_1 + w_2^T w_2 - w_1^T w_2 - w_2^T w_1 = \|w_1 - w_2\|_2^2 \geq 0$$

Constrained optimization

$$\min f(w)$$

$$h_i(w) = 0 \quad i=1, \dots, l$$

$$g_i(w) \leq 0 \quad i=1, \dots, k$$

Primal Problem

p^* = optimal solution

$$L(w, \alpha, \beta) = f(w) + \sum_{i=1}^k \alpha_i g_i(w) + \sum_{i=1}^l \beta_i h_i(w)$$

coefficients must be ≥ 0

coefficients unconstrained

DEF

$$\Theta_p(w) = \max_{\substack{\alpha_i \geq 0 \\ \beta_i}} L(w, \alpha, \beta)$$

If w is feasible (it satisfies all the constraints)
 g_i are all negative and h_i are all 0.
 So we take $\alpha_i = 0$.

If w is not feasible, we take it to infinity through α 's and β 's

$$\Theta_p(w) = \begin{cases} f(w) & w \text{ is feasible} \\ \infty & \text{otherwise} \end{cases}$$

$$\min_w \max_{\substack{\alpha_i \geq 0 \\ b \in B_c}} L(w, \alpha, \beta)$$

To solve above pick w first ^{to minimize} and then solve max.

We must pick a feasible w .

$$\min_w \Theta_p(w) = \text{optimal value of the primal problem} = p^*$$

$$\max_{\substack{\alpha_i \geq 0 \\ b \in B_c}} \min_w L(w, \alpha, \beta) = \max_{\substack{\alpha_i \geq 0 \\ b \in B_c}} \Theta_D(\alpha, \beta)$$

DEF.

$$\Theta_D(\alpha, \beta) = \min_w L(w, \alpha, \beta)$$

d^* = optimal solution

We can show that $d^* \leq p^*$

$$\max_x \min_y \phi(x, y) \leq \min_y \max_x \phi(x, y)$$

$$\min_y \phi(x^*, y) \leq \phi(x^*, y^*) \leq \max_x \phi(x, y^*)$$

JL

Thm.

If f, g_i are convex, g_i strictly feasible, h_i are affine, then

$$p^* = d^*$$

$$h_i(w) = \theta_i^T w + z$$

Strictly feasible means that there is a solution where $g_i(w) < 0$
 $\uparrow$
 strictly $<$

In fact, there exist w^*, α^*, β^* such that

$$p^* = d^* = L(w^*, \alpha^*, \beta^*)$$

$$\frac{\partial}{\partial w_i} L(w^*, \alpha^*, \beta^*) = 0$$

$$\frac{\partial}{\partial \beta_i} L(w^*, \alpha^*, \beta^*) = 0$$

$$\boxed{\alpha_i g_i(w^*) = 0} \quad \forall i$$

$$g_i(w^*) \leq 0$$

$$\alpha_i^* \geq 0$$

$$h_i(w^*) = 0$$

We go back to original Lagrangian, we change $\min \|W\|$ to $\min \frac{1}{2} \|W'\|$

for convenience

$$L(W, b, \alpha) = \frac{1}{2} W^T W + \sum_{i=1}^n \alpha_i (1 - (W^T x^{(i)} + b) y^{(i)})$$

$$\frac{\partial L}{\partial W} = W - \sum_{i=1}^n \alpha_i x^{(i)} y^{(i)} = 0$$

$$\frac{\partial L}{\partial b} = \sum_{i=1}^n \alpha_i y^{(i)} = 0$$

Immediate solution is

$$W = \sum_{i=1}^n \alpha_i x^{(i)} y^{(i)}$$

We plug this W back into the Lagrangian. Because of the nature of

$$\frac{\partial L}{\partial b} = 0,$$

b will not figure into the equation

$$\Theta_0(\alpha) = \sum_{i=1}^n \alpha_i - \frac{1}{2} \sum_{i,j} y^{(i)} y^{(j)} \alpha_i \alpha_j \underline{x^{(i)T} x^{(j)}}$$

$$\alpha_i \geq 0$$

This inner product can be replaced by a kernel.

We can compute sol'n of the dual (even use kernels).

Now we have $\alpha_1, \dots, \alpha_n$. What is the classifier?

$$x \rightarrow \text{sign}(W^T x + b)$$

$$\text{sign}\left(\sum_{i=1}^n \alpha_i y^{(i)} (x^T x^{(i)}) + b\right)$$

$$b = \left(\max_{\substack{i \\ y^{(i)} = -1}} w^T x^{(i)} - \min_{\substack{i \\ y^{(i)} = +1}} w^T x^{(i)} \right) / 2$$

We may allow for an error, i.e.

$$(w'^T x^{(i)} + b) y^{(i)} \geq 1 - \epsilon, \quad i \in \{1, \dots, n\}$$

$$\min \frac{1}{2} \|w'\|_2^2 + \sum \epsilon$$

2.10.2009

Last Class

We were searching for maximum-margin classifiers in form for which we can use kernel trick.

$$\min \frac{1}{2} \|w\|_2^2$$

$$y^{(i)}(w^T x^{(i)} + b) \geq 1$$

Dual:

$$\max \sum \alpha_i - \frac{1}{2} \sum y^{(i)} y^{(j)} \alpha_i \alpha_j (x^{(i)})^T (x^{(j)})$$

$$\alpha_i \geq 0$$

$$\sum \alpha_i y^{(i)} = 0$$

↑
a constraint
for each
 $x^{(i)}, y^{(i)}$
data point

We want to add constraints to allow violations

Modification: (Add slack variables)

$$\min \frac{1}{2} \|w\|_2^2 + C \sum_{i=1}^D \epsilon_i \quad \epsilon_i \geq 0$$

$$y^{(i)}(w^T x^{(i)} + b) \geq 1 - \epsilon_i$$

Dual:

Same as above with $0 \leq \alpha_i \leq C$

For the project, find the optimal solution for the dual problem above.

Sequential Minimal Optimization (SMO)

Pick two indices i, j

Try to modify α_i, α_j to

- increase objective
- stay in the feasible region

$$\alpha_i = z$$

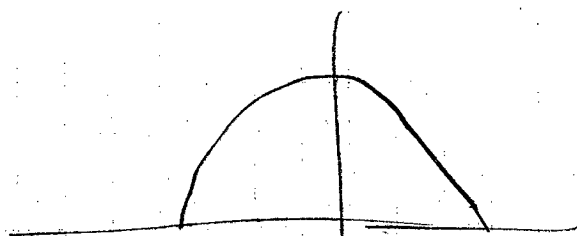
$$\alpha_i y^{(i)} + \alpha_j y^{(j)} = - \sum_{k \neq i, j} \alpha_k y^{(k)} = A$$

$$\alpha_j = (A - y^{(i)} z) y^{(j)}$$

In the end, we are optimizing a quadratic function on an interval.

$$\max a_2 z^2 + a_1 z + a_0 \text{ on } [c_1, c_2]$$

If $a_2 < 0$



$$z z a_2 + a_1 = 0 \rightarrow z = \frac{-a_1}{2a_2}$$

output is
$$\begin{cases} f\left(\frac{-a_1}{2a_2}\right), & \text{if } \left(\frac{-a_1}{2a_2}\right) \text{ is in } [c_1, c_2] \\ \max\{f(c_1), f(c_2)\}, & \text{otherwise} \end{cases}$$

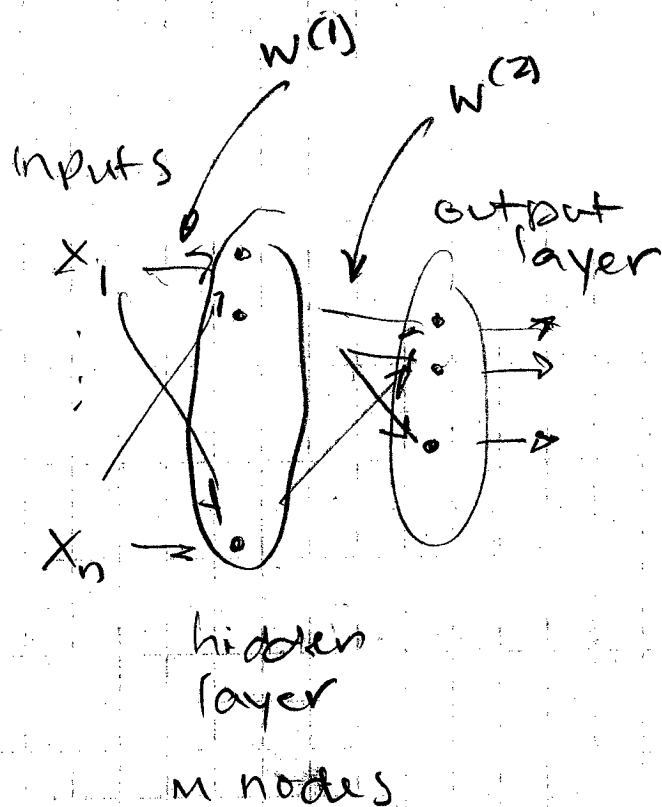
"Suggested initial value of α_i is 0, but it shouldn't matter".

Neural Networks

$$\begin{matrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{matrix} \rightarrow \sum_{i=1}^n w_i x_i$$

perceptron

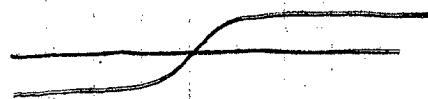
A neural network is a "network of perceptrons."



We have to add non-linearity to get something interesting.

We use $h()$.

For example
$$h(x) = \frac{-e^{-x} + e^x}{e^{-x} + e^x}$$

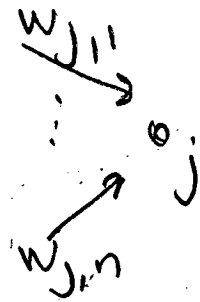


lets use $h()$ for the hidden layer and $\sigma()$ for the output layer.

If the perceptrons have weights

$w_{ji}^{(1)}$ for the first layer

$w_{kj}^{(2)}$ for the second layer



The function computed is

$$y_k = \sigma\left(\sum_j w_{kj}^{(2)} h\left(\sum_i w_{ji}^{(1)} x_i\right)\right)$$

The goal is to find good w using $\bar{x}^{(1)}, \dots, \bar{x}^{(D)}$ and $\bar{t}^{(1)}, \dots, \bar{t}^{(D)}$.

Find w 's such that

$$\sum_{l=1}^D \sum_{k=1}^K (y_k^{(l)} - t_k^{(l)})^2 \text{ is minimized.}$$

$$y_k^{(l)} = \sigma\left(\sum_j w_{kj}^{(2)} h\left(\sum_i w_{ji}^{(1)} x_i^{(l)}\right)\right)$$

To find w , we use gradient descent (back propagation)

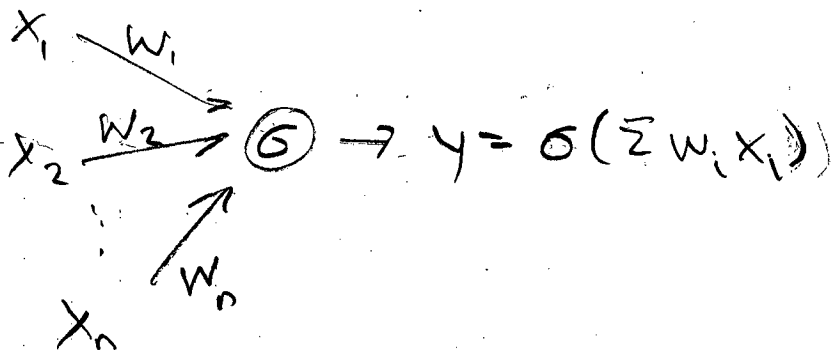
$$E(w, x, \tau)$$

$$\nabla w = \frac{\partial E}{\partial w}$$

learning rate

$$\text{Let } w = w - \eta \cdot \nabla w$$

Let's start by solving for 1 perceptron.



$$E(w, x, T) = (y - T)^2$$

What is $\frac{\partial E}{\partial w} = \left(\frac{\partial E}{\partial w_1}, \dots, \frac{\partial E}{\partial w_n} \right)$

$$\frac{\partial E}{\partial w_l} = 2(y - T) \sigma'(\sum w_i x_i) \cdot x_l$$

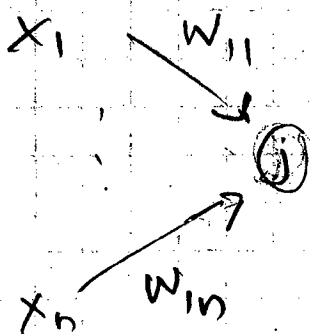
Chain Rule

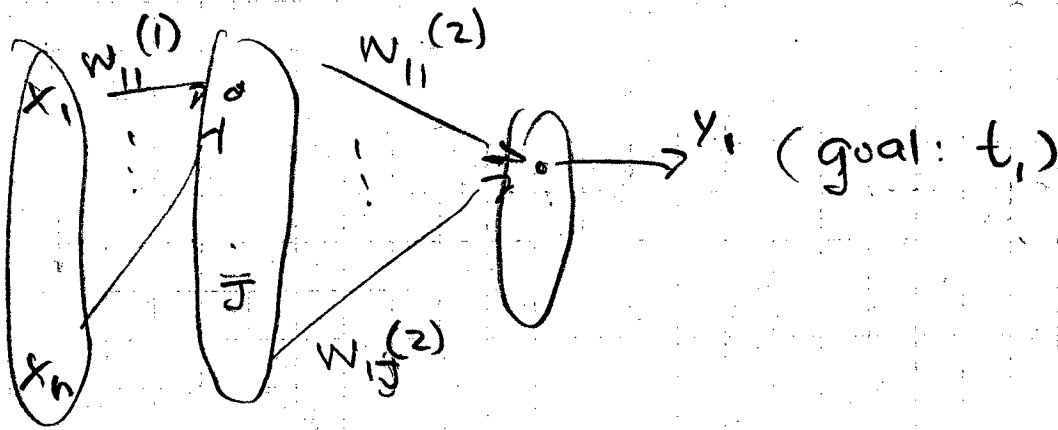
$f(x_1, \dots, x_n)$, x_i is a function of t

$$\frac{df}{dt} = \sum_{i=1}^n \frac{df}{dx_i} \frac{dx_i}{dt}$$

Let the activation of unit j be

$$a_j = \sum w_{ji} x_i$$





$$E = \sum_{k=1}^K \sum (y_k - t_k)^2$$

$$\frac{\partial E}{\partial w_{11}^{(2)}} = 2(y_1 - t_1) \frac{\partial y_1}{\partial w_{11}^{(2)}}$$

$$y_1 = \sigma\left(\sum_{j=1}^J w_{1j}^{(2)} z_j\right)$$

$$\frac{\partial y_1}{\partial w_{11}^{(2)}} = \sigma'(\dots) \cdot z_1$$

For the first layer

$$E = \sum_{k=1}^K (y_k - t_k)^2, \quad y_k = \sigma(\sum w^{(2)} h(\sum w^{(1)} \dots))$$

$$\frac{\partial E}{\partial w_{11}^{(1)}} = \sum \frac{\partial E}{\partial a_k} \cdot \frac{\partial a_k}{\partial w_{11}^{(1)}}$$

2.11.2009

Tutorial

$$(1.31) \quad H(X) + H(Y) \geq H(X, Y)$$

$$H(X) = - \int p(x) \log p(x) dx = - \iint p(x, y) \log p(x) dx dy$$

$$H(X) + H(Y) - H(X, Y) =$$

$$- \iint p(x, y) \log \frac{p(x)p(y)}{p(x, y)} dx dy =$$

$$E_{x, y} \left(-\log \frac{p(x)p(y)}{p(x, y)} \right) \geq -\log E_{x, y} \frac{p(x)p(y)}{p(x, y)} = 0$$

$$\forall_{x, y} \frac{p(x)p(y)}{p(x, y)} \text{ must be } 1 \text{ for } =$$

$$(1.33) \quad H[Y|X] = 0 \Rightarrow \forall x \ p(x) > 0 \exists! y \text{ s.t. } p(y|x) = 1$$

$$H(Y|X) = - \sum_{x, y} p(x, y) \log p(y|x)$$

$$= - \sum_{\substack{x \\ p(x) > 0}} p(x) \sum_y p(y|x) \log p(y|x) = 0$$

always > 0 always ≤ 0

Proof: Assume $\exists x, y$ s.t. $p(x) > 0$ and $0 < p(y|x) < 1$

But then there is a term
 $p(y|x) \log p(y|x) < 0$

and $H[Y|X] > 0$ contradiction

(1.32)

 $y = Ax$, A nonsingular

$$H[y] = H[x] + \ln |A|$$

$$p_y(y) = p_x(x) \left| \frac{dx}{dy} \right|$$

$$\left| \frac{dy}{dx} \right| = |A|$$

$$p_x(y) = \frac{p_x(x)}{|A|}$$

$$H[y] = - \int p_y(y) \ln p_y(y) dy = - \int \frac{p_x(x)}{|A|} \ln \frac{p_x(x)}{|A|} |A| dx$$

Jacobian
determinant

$$= - \int p_x(x) \ln p_x(x) dx + \int p_x(x) \ln |A| dx$$

$$= H[x] + \ln |A|$$

2.12.2009

Graphical Models

We want to describe a distribution of $x_1, \dots, x_n$ in a compact form.

Example

$$x_i \in \{0, 1\}$$

To describe $x_1, \dots, x_n$ requires 2^n numbers

Perhaps x_1, x_2 and x_3 are generated by some process.

Example

$$x_1 \rightarrow x_2 \rightarrow x_3$$

$$\left\{ \begin{array}{l} P(x_1=0) \\ P(x_1=1) \end{array} \right\} f(x_1)$$

$$\left\{ \begin{array}{l} P(x_2=0|x_1=0) \\ P(x_2=1|x_1=1) \end{array} \right\} f(x_2|x_1)$$

$$\left\{ \begin{array}{l} P(x_3=0|x_2=0) \\ P(x_3=1|x_2=1) \end{array} \right\} f(x_3|x_2)$$

Directed Graphical Model

- directed graph G
acyclic

- vertices correspond to $x_1, \dots, x_n$
(each vertex one variable)

- we have $f(x_v | x_{\pi_v})$ — distribution of x_v given parents
for each vertex v

Let's use notation $x_{\{2,3\}} = (x_2, x_3)$

The distribution of $X_1, \dots, X_n$ described by the model is

$$P(X_1=a_1, X_2=a_2, \dots, X_n=a_n) = \prod_{v \in V} f(X_v=a_v | X_{\pi_v}=a_{\pi_v})$$

Let's go back to the example.

$$X_1 \rightarrow X_2 \rightarrow X_3$$

$$f_1(X_1=0) = 1/2$$

$$f_3(X_3=0 | X_2=0) = 2/3$$

$$f_2(X_2=0 | X_1=0) = 2/3$$

$$f_3(X_3=0 | X_2=1) = 1/3$$

$$f_2(X_2=0 | X_1=1) = 1/3$$

$$\text{Then, } P(0, 1, 0) = P(X_1=0, X_2=1, X_3=0) = \\ 1/2 \cdot 1/3 \cdot 1/3$$

Note that we are using

$$(*) \quad \sum_{a_v} f(X_v=a_v | X_{\pi_v}=a_{\pi_v}) = 1 \text{ for any } a_{\pi_v}$$

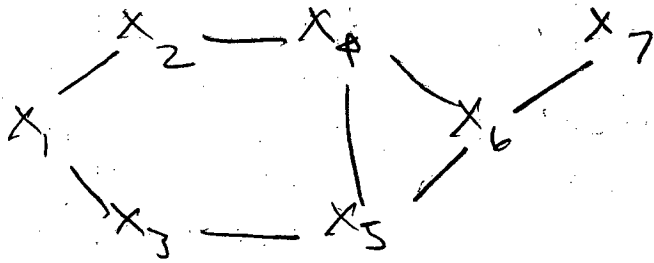
If f 's above satisfy $(*)$, then P is a probability.

Fact 1:

$$\text{Fact 2: } P(X_v | X_{\pi_v}) = f(X_v | X_{\pi_v})$$

Undirected Graphical Model

20



- undirected graph
- $f(x_c)$ for every maximal clique

In the drawing above, there are 6 maximal cliques

A Maximal clique is a clique that is not a subset of another clique.

$$P(x_1, \dots, x_n) \propto \prod_c f(x_c) \\ = \frac{1}{Z} \prod_c f(x_c)$$

$$P(x_1=a_1, \dots, x_7=a_7) = \frac{1}{Z} \underbrace{f_{13}(a_1, a_3) \cdot f_{456}(a_4, a_5, a_6) \cdot \dots}_{6 \text{ terms}}$$

6 terms

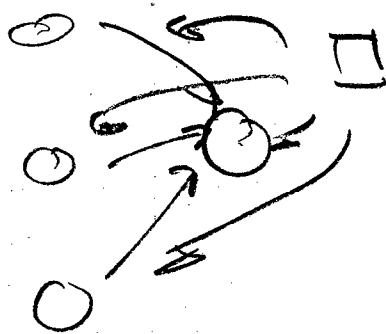
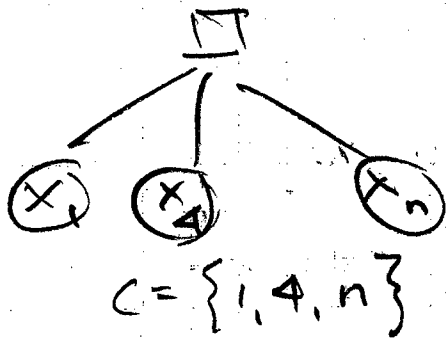
The above holds for the case when we are specifying values for all x_i .

Factor Graphs

collection of subsets $C \subseteq \{1, \dots, n\}$

f_c for each of your chosen C

$$P(x_1, \dots, x_n) \propto \prod f_c(x_c)$$



$$f_{\pi_c}^v = f(x_v | x_{\pi_v})$$

-
- To go from factor graph to undirected graph:
- Take G .
 - For each v , add undirected edges between the parents of v .
 - Forget the direction of the original edges. They become undirected.
-

Returning to undirected graphs ...

$$P(X_3=0) = \frac{1}{2}$$

$$\begin{array}{ccc} 0 & 0 & 0 \\ 0 & 1 & 0 \\ 1 & 0 & 0 \\ 1 & 1 & 0 \end{array}$$

How do we calculate probabilities of subsets of the variables?

$$\begin{aligned} P(X_1=a_1) &= \sum_{a_2} \sum_{a_3} \dots \sum_{a_n} P(X_1=a_1, X_2=a_2, \dots, X_n=a_n) \\ &= \frac{1}{2} \sum_{a_2} \dots \sum_{a_7} f_{12}(a_1, a_2) f_{24}(a_2, a_4) \dots \end{aligned}$$

$$f_{13}(a_1, a_3) f_{35}(a_3, a_5) f_{456}(a_4, a_5, a_6) \cdot f_{67}(a_6, a_7)$$

The above would take D^6 time

What if we precomputed

$$m_{6,7}(a_6) = \sum_{a_7} f_{6,7}(a_6, a_7)$$

for all a_6 , it would take D^2 time.

We have gotten rid of one of summations

$$\frac{1}{2} \sum_{a_2} \sum_{a_3} \sum_{a_4} \sum_{a_5} \sum_{a_6} \dots m_{6,7}(a_6)$$

our summation now takes D^5 time.

Let's try again

$$M_{4,5,6}(a_4, a_5) = \sum_{a_6} f_{4,5,6}(a_4, a_5, a_6) m_{6,7}(a_6)$$

Computing $M_{4,5,6}$ for all a_4, a_5 requires D^3 time.

Now, we have

$$P() = \frac{1}{Z} \sum_{a_2} \sum_{a_3} \sum_{a_4} \sum_{a_5} f_{1,2}(a_1, a_2) f_{1,3}(a_1, a_3) f_{2,4}(a_2, a_4) f_{3,5}(a_3, a_5) \cdot M_{4,5,6}()$$

This requires D^4 time.

Now, what if we tried to eliminate a_3 ?

$$\sum_{a_3} f_{1,3}(a_1, a_3) f_{3,5}(a_3, a_5) = m_{1,5,3}(a_1, a_5)$$

↑
 D^3 time

This leaves

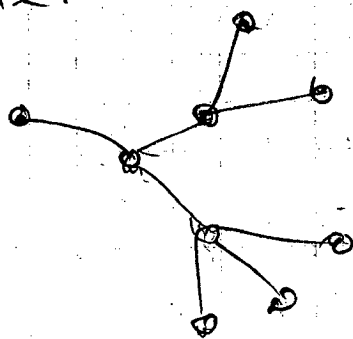
$$\frac{1}{Z} \sum_{a_2} \sum_{a_4} \sum_{a_5} f_{1,2}(a_1, a_2) f_{2,4}(a_2, a_4) m_{1,5,3}(a_1, a_5) \cdot m_{4,5,6}(a_4, a_5)$$

$$Z = \sum_{a_1} \dots \sum_{a_7} \dots$$

↑
 D^4 time

What we have been doing is called the elimination algorithm.

Given a graph G , what order should we choose for elimination? What is the running time?



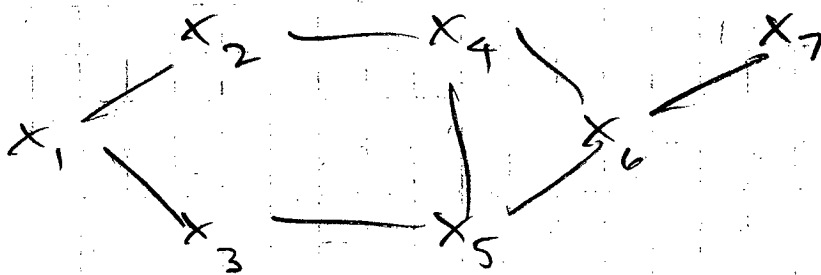
If G is a tree, we use D^2 time to eliminate a leaf node.

If G is a tree, always process a leaf. Running time $\sim nD^2$

In general, the running time is $O^{\text{treewidth}(G)}$

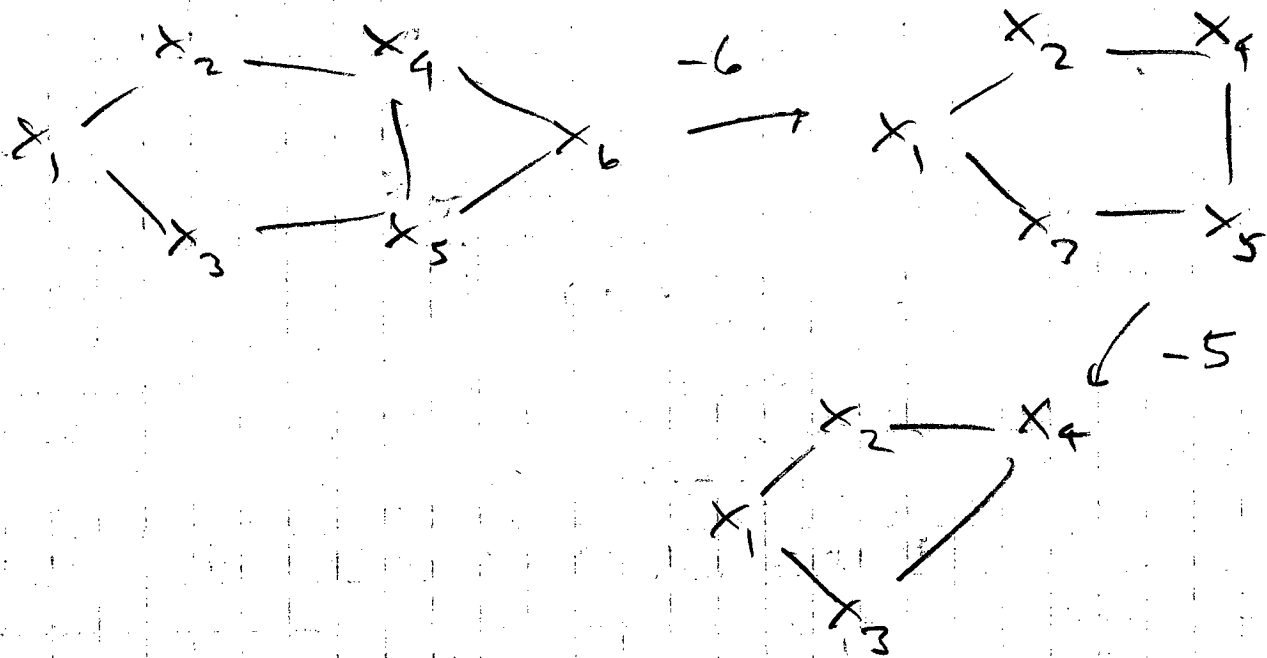
Tree width of G is the minimum, over all possible orderings of the vertices of G , of the maximum degree of an eliminated vertex in the following process:

To eliminate a vertex, put edges between all of its neighbors.



7, 6, 5, 4, 2, 3, 1

Eliminate 7

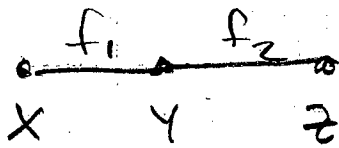


$$tw = 2 + 1$$

tree width of a tree is 2

217 2009

Conditional Independence in Graphical Models

Are X, Z independent?

not necessarily

$$f_1(0,0) = 1$$

$$f_1(1,1) = 1$$

$$f_1(0,1) = 0$$

$$f_1(1,0) = 0$$

$$f_2 = f_1$$

$$P(X=0, Y=1, Z=1) = 0$$

$$P(X=Y=Z=1) = \frac{1}{2}$$

$$P(X=Y=Z=0) = \frac{1}{2}$$

 X, Z are not independentAre X, Z independent given Y ? U, W are independent $\Leftrightarrow$

$$\forall u, w \quad P(U=u \wedge W=w) = P(U=u) P(W=w)$$

 U, W are independent given $\phi \Leftrightarrow$

$$\forall u, w, q \quad P(U=u \wedge W=w \mid \phi=q) =$$

$$P(U=u \mid \phi=q) P(W=w \mid \phi=q)$$

In the concrete example,

$$P(X=1 \wedge Z=0 | Y=0)$$

↪ Check all 8 cases to show that X, Z are independent given Y for this example

In general

$$P(X, Z | Y) \stackrel{?}{=} P(X | Y) P(Z | Y)$$

$$\frac{P(X, Y, Z)}{P(Y)} = \frac{f_1(X, Y) f_2(Y, Z)}{\sum_{X'} \sum_{Z'} f_1(X', Y) f_2(Y, Z')}$$

$$P(X | Y) = \frac{P(X, Y)}{P(Y)} = \frac{\sum_{Z'} f_1(X, Y) f_2(Y, Z')}{P(Y)} =$$

$$\frac{f_1(X, Y) \sum_{Z'} f_2(Y, Z')}{\sum_{X'} \sum_{Z'} f_1(X', Y) f_2(Y, Z')}$$

$$P(X | Y) P(Z | Y) = \frac{f_1(X, Y) \sum_{Z'} f_2(Y, Z') \cdot f_2(Y, Z) \cdot \sum_{X'} f_1(X', Y)}{(\sum_{X'} \sum_{Z'} f_1(X', Y) f_2(Y, Z'))^2}$$

$$\frac{P(X, Y, Z)}{P(Y)} = \frac{f_1(X, Y) f_2(Y, Z)}{(\sum_{X'} f_1(X', Y)) (\sum_{Z'} f_2(Y, Z'))}$$

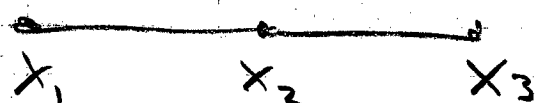
↑
equal
↙

So, it's true that $X \perp\!\!\!\perp Z \mid Y$ (X and Z are independent given Y)

For $S \subseteq V$ — vertices $= v_1, \dots, v_k$

Let $X_S = (X_{v_1}, \dots, X_{v_k})$

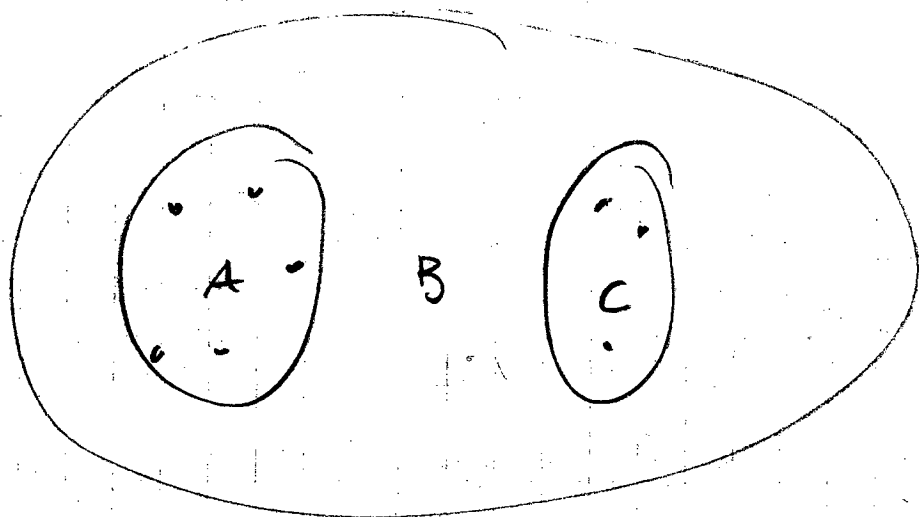
For example



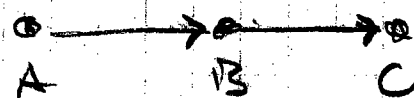
$$X_{\{1,3\}} = (X_1, X_3)$$

$A, B, C \subseteq V$

In general $X_A \perp\!\!\!\perp X_C \mid X_B$ if B separates A from C



Let's look at directed graphs.



Is $X_A \perp\!\!\!\perp X_C \mid X_B$? YES

Is $X_A \perp\!\!\!\perp X_C$? NO

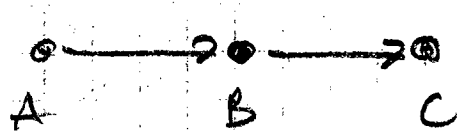
$$P(A, B, C) = P(A) P(B|A) P(C|B)$$

We want to show $P(A, C|B) = P(A|B) P(C|B)$

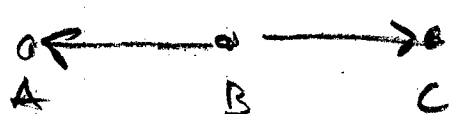
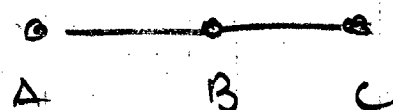
$$P(A, C|B) = \frac{P(A, B, C)}{P(B)} = \frac{P(A) P(B|A) P(C|B)}{P(B)} = P(A|B) P(C|B)$$

$P(A|B)$
By Bayes

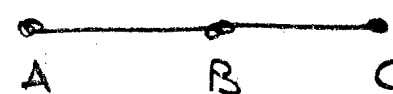
Note that



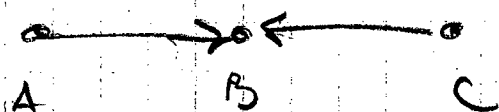
reduces
to



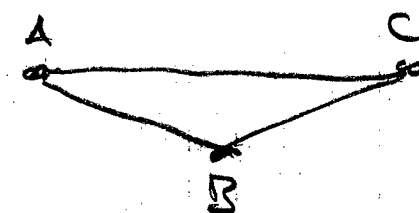
reduces
to



what about



reduces
to



$$X_A \perp\!\!\!\perp X_C \mid X_B \quad ? \quad \text{NO}$$

$$X_A \perp\!\!\!\perp X_C \quad ? \quad \text{YES}$$

$$P(A, B, C) = P(A) P(C) P(B|A, C)$$

$$\frac{P(A, B, C)}{P(B|A, C)} = P(A, C) = P(A) P(C)$$

$$X_A \in_R \{0, 1\}$$

$$X_C \in_R \{0, 1\}$$

$$X_B = X_A + X_C \pmod{2}$$

$$\text{Let } X_B = 0$$

$$\text{Then } P(X_A = 0 \wedge X_C = 1 \mid X_B = 0) \stackrel{?}{=}$$

$$\stackrel{?}{=} 0$$

$$P(X_A = 0 \mid X_B = 0) \cdot P(X_C = 1 \mid X_B = 0)$$

$$\stackrel{?}{=} \frac{1}{2}$$

$$\stackrel{?}{=} \frac{1}{2}$$

$$\text{So } X_A \not\perp X_C \mid X_B$$

To prove independence must exist for all edge functions

To disprove $\exists$ edge function s.t. not independent.

In undirected, conditional independence implies that there is no path from A to C that does not go through B .

For directed,

We say that A, C are d -connected w.r.t. B if there exists an undirected path from $a \in A$ to $c \in C$ such that every vertex $v \in P$ ^{for}

1. if v is non-blocking, v is not in B

2. if v is blocking, v or its descendant is in B

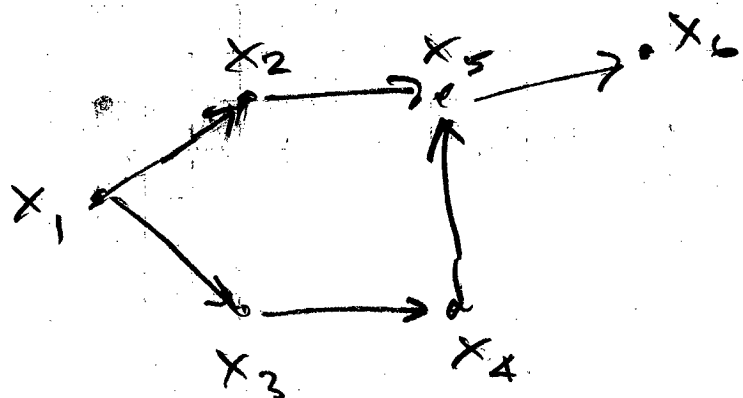
Blocking $\rightarrow \bigcirc \leftarrow$
 $\downarrow$
 V

DEF
 A, C are d-separated by B if they are not d-connected w.r.t. B

THM

If A, C are d-separated by B , then

$$X_A \perp\!\!\!\perp X_C \mid X_B$$

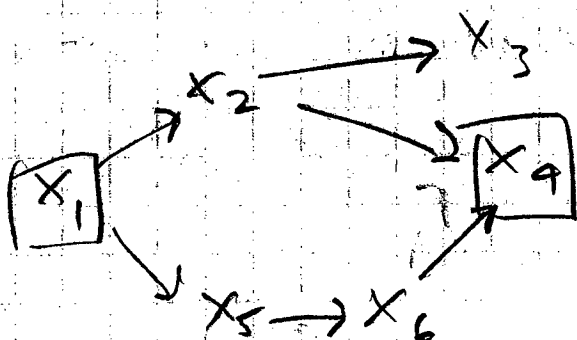


$$X_1 \perp\!\!\!\perp X_5 \mid X_2, X_3 \quad ? \quad \text{YES}$$

Look at paths

X_1, X_2, X_5 separated

X_1, X_3, X_4, X_5 separated



$$X_2 \perp\!\!\!\perp X_5 \mid X_1, X_4 \quad ?$$

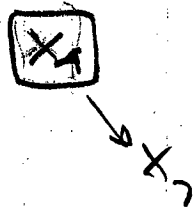
not
 X_5, X_1, X_2 is a d-connecting path

So, X_2, X_5 is d-separated by X_1 .

X_5, X_6, X_4, X_2 is a d-connecting path.

$\therefore X_2 \perp\!\!\!\perp X_5 \mid X_1, X_4$ is FALSE

If you added X_7 to the graph st.

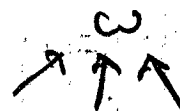


then $X_2 \perp\!\!\!\perp X_5 \mid X_1, X_7$ is also FALSE. This shows the descendent case.

Proof of THM ^{directed} (G is _{acyclic} graph (DAG))

Induction on the number of vertices in G .
Skip the base, single vertex, case

Let w be a sink of G



Case: $w \notin A \cup B \cup C$

A and C are d-separated by B in G

$\Downarrow$

A and C are d-separated by B in $G-w$

$\Downarrow$

$X_A \perp\!\!\!\perp X_C \mid X_B$

We choose a sink w because there will always be a sink in a directed graph.

2.19.2009

Continuation of D-Separation Proof...Directed graphical model $G=(V,E)$ $A, B, C \subseteq V$ (disjoint)DEF A is d-connected to C w.r.t. B if
 $\exists$ directed path $P = a_1, \dots, C$ $a \in A$ $C \in C$

- For every $V \in P$
- If V is collider then V or an ancestor of V is in B
 - If V is not a collider, then $V \notin B$

 A is d-separated from C by B if A is not d-connected to C w.r.t. B TheoremIf A is d-separated from C by B , then

$$X_A \perp\!\!\!\perp X_C \mid X_B$$

ProofInduction on $|V|$

- Let w be a vertex of out-degree 0 in G (i.e. w is a sink)
- Let $V' = V - w$, $G' = (V', E')$

Fact:

 $X_{V'}$ factorizes over G'

Case 1: $w \notin A \cup B \cup C$

Clearly A is d -separated from C by B in G'

By IH (inductive hypothesis)

$$X_A \perp\!\!\!\perp X_C \mid X_B \text{ in } G'$$

$$\downarrow$$

$$X_A \perp\!\!\!\perp X_C \mid X_B \text{ in } G$$

Case 2: $w \in A$

• Let $A' = A - w$

• Let P be the set of parents of w which are not in B

Claim 1:

A' is d -separated from C by B in G'

Claim 2: P is d -separated from C by B in G'

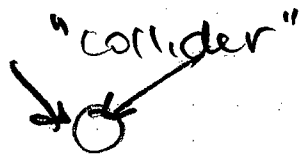
Claim 1 is satisfied by assumption of theorem

Proof of Claim 2:

Suppose that the claim is false, i.e.
 $\exists$ good path from $c \in C$ to $p \in \text{Parents}(w) - B$



Then the path extended by (p, w) is good, too



If this was true, then A is d -connected to C w.r.t. B . This is a contradiction.

By IH,

$$X_{A' \cup P} \perp\!\!\!\perp X_C \mid X_B$$

$$X_P \perp\!\!\!\perp X_C \mid X_B$$

Fact

$$X_W \perp\!\!\!\perp X_{\text{any}} \mid X_P$$

anything that
doesn't include
w or p

$$A \perp\!\!\!\perp B \mid C \cup D, A \perp\!\!\!\perp D \mid C \Rightarrow A \perp\!\!\!\perp B \cup D \mid C \quad (2)$$

$$U \perp\!\!\!\perp V \mid W \stackrel{\text{def}}{\Leftrightarrow} P(U|V, W) = P(U|W)$$

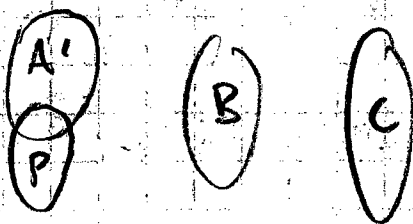
Given the above

$$\left. \begin{aligned} P(A|B, C, D) &= P(A|C, D) \\ P(A|C, D) &= P(A|C) \end{aligned} \right\} P(A|B, C, D) = P(A|C)$$

$A \perp\!\!\!\perp B \cup D \mid C$ ←

Going back to our fact

$$X_W \perp\!\!\!\perp X_C \mid X_{A' \cup P \cup B}$$



We want to conclude that

$$X_{A' \cup \{w\}} \perp\!\!\!\perp X_C \mid X_B$$

Repeating, ...

By the IH

$$X_{A' \cup P} \perp\!\!\!\perp X_C \mid X_B \quad (2)$$

$$X_w \perp\!\!\!\perp X_C \mid X_{B \cup A' \cup P} \quad (1)$$

Conclusion

$$X_{A' \cup w \cup P} \perp\!\!\!\perp X_C \mid X_B \Rightarrow X_{A' \cup w} \perp\!\!\!\perp X_C \mid X_B$$

Map to laws on previous page

A	X_C
B	X_w
C	X_B
D	$X_{A' \cup P}$

$$A \perp\!\!\!\perp B \mid C \Leftrightarrow P(A \mid B, C) = P(A \mid C) \quad D1$$

$$B \perp\!\!\!\perp A \mid C \Leftrightarrow P(B \mid A, C) = P(B \mid C)$$

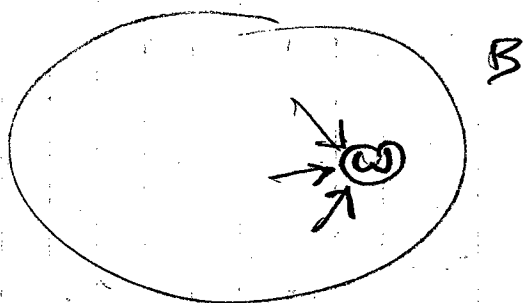
$$A \perp\!\!\!\perp B \mid C \Leftrightarrow P(A \mid C) P(B \mid C) = P(A, B \mid C) \quad D2$$

$$P(A \mid C) P(B \mid C) = P(A, B \mid C) = P(A \mid B, C) P(B \mid C)$$

$\therefore$ Independence is symmetric

Case: $w \in B$

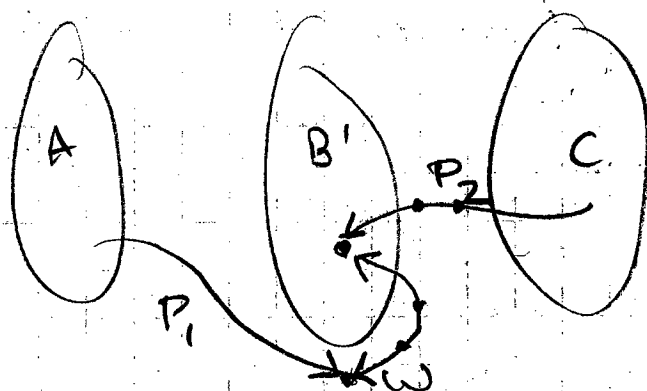
Claim
 A is separated from c by $B' = B - w$



Proof

Let P be a path d -connecting A to c w.r.t. B' . Then $w \notin P$ and w is a path d -connecting A to c w.r.t. B . This is a contradiction.

Claim
 w is d -separated from A or c by B'



Suppose the claim is false. Then $P_1 + P_2$ is a path d -connecting A, c w.r.t. B'_2 .

Say, w.o.l.o.g. that w is d -separated from c by B' . Then

$A \cup w$ is d -separated from c by B'

By IH

$$X_{A \cup W} \perp\!\!\!\perp X_C \mid X_{B'}$$

$$A \perp\!\!\!\perp (B \cup D) \mid C \Rightarrow A \perp\!\!\!\perp B \mid C \cup D \quad \text{Rule}$$

Using the above

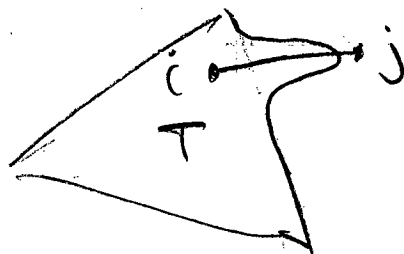
$$X_A \perp\!\!\!\perp X_C \mid X_{B' \cup W} \Rightarrow X_A \perp\!\!\!\perp X_C \mid X_B$$

Sum-Product Algorithm

- Assume an undirected graphical model where G is a tree
 - We have $f_{ij}(x_i, x_j)$ for every edge $(i, j) \in E$
- $$P(x_1, \dots, x_n) \propto \prod_{i,j} f_{ij}(x_i, x_j)$$

Goal: Compute $P(x_i)$ m_{ij} is a message from i to j

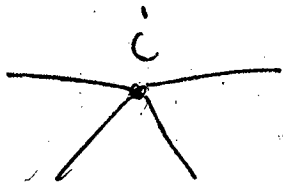
$$x_1, \dots, x_n \in \{1, \dots, D\}$$

 $m_{ij}(x_j) \Leftarrow$ "weight of the subtree rooted at edge i, j if value of x_j is given"


$$m_{ij}(x_j) = \sum_{x_i} f_{ij}(x_i, x_j) \prod_k m_{ki}(x_i)$$

neighbors of
i excluding j

Messages start at the leaves and proceed toward the root.



$$P(x_i) = \frac{1}{Z} \prod_k m_{ki}(x_i)$$

neighbors
of i

$$Z = \sum_{x_i} \prod_k m_{ki}(x_i)$$

2.24.2009

EM Algorithm

How do we implement ML?

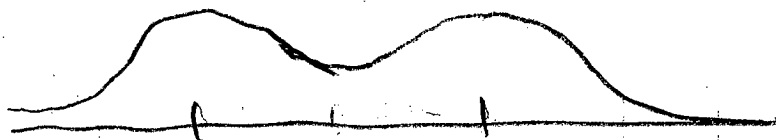
Let's use an example model of mixtures of Gaussians. In one dimension

$$f(x) = \sum_{i=1}^k \pi_i N(x | \mu_i, \Sigma_i) \quad ; \quad \Sigma_i = \frac{1}{6_i^2}$$

(mixture coefficients $\sum \pi_i = 1$)

Input: datapoints $x_1, \dots, x_n \in \mathbb{R}$ output

$$\frac{1}{2} N(x | -1, 1) + \frac{1}{2} N(x | +1, 1)$$



Setting of $\pi_1, \dots, \pi_k, \mu_1, \dots, \mu_k, \Sigma_1, \dots, \Sigma_k$ maximizing the likelihood of the input data $(x_1, \dots, x_n)$.

What we are maximizing is

$$P(x_1, \dots, x_n \mid \begin{matrix} \pi_1, \dots, \pi_k \\ \mu_1, \dots, \mu_k \\ \Sigma_1, \dots, \Sigma_k \end{matrix})$$

Since we don't like products, we take the logarithm of the likelihood. E2

$$\begin{aligned}\log P(\dots) &= \sum_{i=1}^n \log P(x_i | \dots) = \\ &= \sum_{i=1}^n \log \sum_{j=1}^k \pi_j e^{\frac{-(x_i - \mu_j)^2 / 2\sigma_j^2}{\sqrt{2\pi} \sigma_j}} = \Phi\end{aligned}$$

Let's think of the mixture model as a model with a hidden variable ("that tells us which Gaussian to use")

The hidden variables are

$$z_1, \dots, z_k, \quad \sum z_i = 1, \quad z_i \in \{0, 1\}$$

binary variable

$$P(z_i = 1) = \pi_i$$

$$P(x | z_i = 1) = N(x | \mu_i, \sigma_i)$$

This gives the same distribution as f .
For each x , only one of the z_i is 1.

Suppose we knew the parameters π_i , μ_i , and σ_i , which component did x come from?

$$\begin{aligned}P(z_j = 1 | x) &= \frac{P(x | z_j = 1) P(z_j = 1)}{P(x)} \\ &= \frac{\pi_j N(x | \mu_j, \sigma_j)}{\sum_{l=1}^k \pi_l N(x | \mu_l, \sigma_l)}\end{aligned}$$

$$\frac{\partial \varphi}{\partial \mu_j} = \sum_{i=1}^n \frac{\partial \varphi_i}{\partial \mu_j}$$

$$\text{Let } \varphi_i = \log \sum_{j=1}^k \pi_j N(x_i | \mu_j, \Sigma_j)$$

$$\text{Since } (\log f)' = \frac{f'}{f},$$

$$\frac{\partial \varphi_i}{\partial \mu_j} = \frac{\frac{\partial N(x_i | \mu_j, \Sigma_j)}{\partial \mu_j} \pi_j}{\sum_{l=1}^k \pi_l N(x_i | \mu_l, \Sigma_l)}$$

$$\frac{\partial N(x | \mu_j, \Sigma_j)}{\partial \mu_j} = N(x | \mu_j, \Sigma_j) \Sigma_j^{-1} (x - \mu_j)$$

So,

$$\frac{\partial \varphi_i}{\partial \mu_j} = \frac{N(x_i | \mu_j, \Sigma_j) \cdot \Sigma_j^{-1} (x_i - \mu_j) \pi_j}{\sum_{l=1}^k \pi_l N(x_i | \mu_l, \Sigma_l)}$$

and

$$\frac{\partial \varphi}{\partial \mu_j} = \sum_{i=1}^n \frac{\pi_j (x_i - \mu_j) \Sigma_j^{-1} N(x_i | \mu_j, \Sigma_j)}{\sum_{l=1}^k \pi_l N(x_i | \mu_l, \Sigma_l)} = 0$$

$$= \sum_{i=1}^n \delta_{ij} (x_i - \mu_j) \Sigma_j^{-1} = 0$$

$$\delta_{ij} = \frac{\pi_j N(x_i | \mu_j, \Sigma_j)}{\sum_{l=1}^k \pi_l N(x_i | \mu_l, \Sigma_l)} = P(z_j = 1 | x_i)$$

Now, we can solve for all μ_j 's.

$$\mu_j = \frac{1}{N_j} \sum_{i=1}^n \delta_{ij} x_i, \quad N_j = \sum_{i=1}^n \delta_{ij} \quad (1)$$

If we know δ_{ij} 's, then we can compute μ_j .

We will not derive the expression for $\bar{z}_j, \pi_j$.

$$\Sigma_j = \frac{1}{N_j} \sum_{i=1}^n \delta_{ij} (x_i - \mu_j)^2, \quad \text{where the } \mu_j \text{ are calculated by previous formula (1)} \quad (2)$$

$$\pi_j = \frac{N_j}{\sum_{k=1}^K N_k} \quad (3)$$

EM Algorithm

$$\pi^{(0)}, \mu^{(0)}, \bar{z}^{(0)}$$

- Initialize π_i 's, μ_i 's, and $\bar{z}_i$'s somehow
- E-step: Compute $\delta_{ij}^{(t+1)}$ using $\pi^{(t)}, \mu^{(t)}, \bar{z}^{(t)}$
- M-step: Compute $\pi^{(t+1)}, \mu^{(t+1)}, \bar{z}^{(t+1)}$ using (1), (2), (3), and $\delta_{ij}^{(t+1)}$

EM Algorithm (In general)

Model: X observable

Z hidden (latent) variables

θ parameters

Goal: Find settings of θ maximizing

$$P(X|\theta)$$

$$P(X, Z|\theta)$$

$$P(X|\theta) = \sum_z P(X, z|\theta)$$

Goal:

$$\arg \max_{\theta} \log \sum_z P(X, z|\theta)$$

$$L(q, \theta) := \sum_z q(z|x) \log \frac{P(X, z|\theta)}{q(z|x)}$$

EM Algorithm

• Set $\theta^{(0)}$ somehow

• E-step: $q^{(t+1)} = \arg \max_q L(q, \theta^{(t)})$

• M-step: $\theta^{(t+1)} = \arg \max_{\theta} L(q^{(t+1)}, \theta)$

The q 's should correspond to the t_i in the Gaussian case

What should q be in the E-step?

2.26.2009

EM Algorithm for models with hidden variables.

$$\operatorname{argmax}_{\theta} p(x|\theta) \leftarrow \text{GOAL}$$

What we have is $p(x, z|\theta)$

So, the goal is really a marginal distribution

$$\operatorname{argmax}_{\theta} \sum_z p(x, z|\theta)$$

$$L(q, x, \theta) = \sum q(z|x) \log \frac{p(z, x|\theta)}{q(z|x)}$$

$$\text{Let } \lambda(x, \theta) = \log p(x|\theta)$$

EM Algorithm Steps

0. First set θ to something (i.e. $\theta^{(0)}$)

$$1. q^{(t+1)} = \operatorname{argmax}_q L(q, x, \theta^{(t)})$$

$$2. \theta^{(t+1)} = \operatorname{argmax}_{\theta} L(q^{(t+1)}, x, \theta)$$

" q is a distribution on the hidden variable"

The EM algorithm is an unsupervised learning method.

For the mixture of Gaussians, we had

1. calculate $p(z_i | x_j) = \delta_{ij}$
2. Use the above to estimate the parameters of the Gaussians (i.e. μ_i 's, σ_i 's, α_i 's)

In general, it is true that

$$q^{(t+1)} = p(z|x, \theta^{(t)})$$

Let's look at

$$\begin{aligned} \ell(x, \theta) - L(q, x, \theta) &= \\ \sum_z q(z|x) \left(\log p(x|\theta) - \log \frac{p(z, x|\theta)}{q(z|x)} \right) &= \\ \underbrace{\sum_z q(z|x) \log p(x|\theta)}_{\ell(x, \theta)} & \end{aligned}$$

$$\sum_z q(z|x) \left(\log \frac{q(z|x)}{p(z|x, \theta)} \right) = \leftarrow \text{This is the KL difference between } q \text{ and } p(z|x, \theta)$$

$$KL(q(z|x) \parallel p(z|x, \theta))$$

$$\left(\frac{p(z, x|\theta)}{p(x|\theta)} = p(z|x, \theta) \right)$$

$KL(\cdot)$ is always positive. A minimum occurs when 2 distributions are equal.

$$\ell(x, \theta) \geq L(q, x, \theta) \text{ since } KL(\dots) \geq 0$$

Fact 1 $\uparrow$

Fact 2 $\downarrow$

$$\underset{q}{\operatorname{argmax}} L(q, x, \theta) = p(z|x, \theta)$$

Fact 2 is true because maximizing $L(q, x, \theta)$ is the same as minimizing $\ell(x, \theta) - L(q, x, \theta) = KL(\dots)$

$$\text{Is } \ell(x, \theta^{(t+1)}) \geq \ell(x, \theta^{(t)}) ?$$

In step 2, $L(q^{(t+1)}, x, \theta)$ will be at least

$$L(q^{(t+1)}, x, \theta) \geq \ell(x, \theta^{(t)})$$

$$\therefore \ell(x, \theta^{(t+1)}) \geq \ell(x, \theta^{(t)})$$

After the E-step

$$L(q^{(t+1)}, x, \theta^{(t)}) = \ell(x, \theta^{(t)})$$

After the M-step

$$L(q^{(t+1)}, x, \theta^{(t+1)}) \geq L(q^{(t+1)}, x, \theta^{(t)})$$

and

$$\ell(x, \theta^{(t+1)}) \geq L(q^{(t+1)}, x, \theta^{(t+1)})$$

$$\therefore \ell(x, \theta^{(t+1)}) \geq \ell(x, \theta^{(t)})$$

Sampling

We have a density function p ? How do we sample x with density p ?

Approach #1:

Sample $x \sim \text{uniform}[0,1]$

Output $z = f(x)$

g is the inverse of f

$$P(Z \leq z) = P(f(x) \leq z) = P(x \leq g(z)) = g(z)$$

because of uniform distribution

$$g(x) = \inf \{ u ; f(u) \geq x \}$$

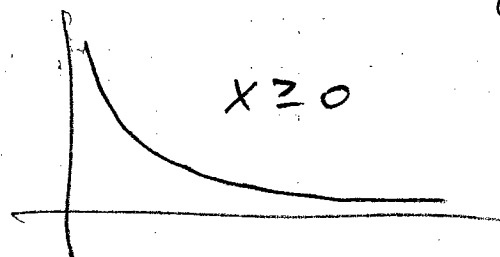
Given density p , compute the cumulative distribution function

$$g(x) = \int_{-\infty}^x p(t) dt$$

Take f to be the inverse of g .

Example ($\text{Exp}(x)$) - Sample from exponential distribution

$$p(x) = \lambda e^{-\lambda x}$$



$$g(x) = \int_0^x \lambda e^{-\lambda t} dt = -e^{-\lambda t} \Big|_0^x = 1 - e^{-\lambda x}$$

$$f(1 - e^{-\lambda x}) = x$$

$$u \mapsto \frac{\ln(1-u)}{\lambda}$$

How would we sample from the normal distribution?

$$X \sim \text{Uniform}([0, 1])$$

$$Y \sim \text{Exp}(1)$$

$$Z_1 = \sqrt{2Y} \sin 2\pi X$$

$$Z_2 = \sqrt{2Y} \cos 2\pi X$$

Then Z_1, Z_2 are independent samples from $N(0, 1)$

$$X \sim N(\mu, \sigma)$$

$$\frac{X - \mu}{\sigma} \sim N(0, 1)$$

$$X \sim N(0, 1)$$

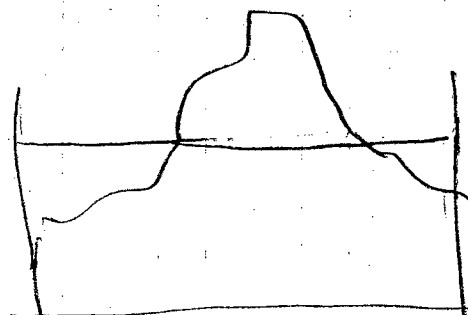
$$\sigma X + \mu \sim N(\mu, \sigma)$$

Rejection Sampling

We have a density function p that we want to sample from

Find q such that

- we know how to sample from q
- $q(x) \geq \frac{1}{k} p(x)$ for reasonably small k



1. Sample x from q

2. Accept x with probability $\frac{p(x)}{k q(x)}$
otherwise

What is the probability of accepting x on the first try?

$$q(x) \cdot \frac{p(x)}{q(x) \cdot k} = \frac{p(x)}{k}$$

$$P(\text{outputting } x \mid \text{accepting on first try}) = p(x)$$

$$P(\text{accepting some } x) = \sum_x q(x) \cdot \frac{p(x)}{q(x)k} = \frac{1}{k}$$

$$p(x) = (2 + \sin x) e^{-x} \cdot \frac{1}{z}$$

z is chosen so that $\int p(x) dx = 1$

Chose $q(x) = e^{-x}$

$$k = \max_x \frac{p(x)}{q(x)} = \max_x (2 + \sin x) \cdot \frac{1}{z} \leq 3$$

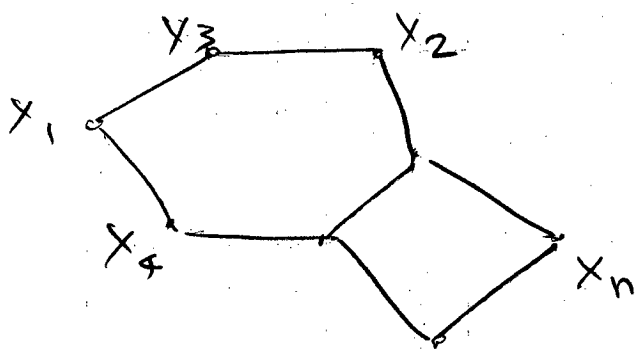
Then, sample $\sim \text{Exp}(1)$.
Accept with probability $\frac{2 + \sin x}{3z}$

Markov Chain Monte Carlo

Example:

Sampling from a distribution with an undirected graphical model.

Gibbs sampler



Pick $i \in \{1, \dots, n\}$

Change

$$x_i \leftarrow p(x_i \mid x_1, \dots, x_{i-1}, x_{i+1}, \dots, x_n)$$

We want to sample from an "ugly" distribution.
We design a process.

A Markov chain is a sequence of random variables

$$X_1, \dots, X_n, \dots$$

such that

Let's assume these are discrete.

$$P(X_{n+1} | X_n) = P(X_{n+1} | X_1, \dots, X_n)$$

This is time homogeneous MC if

$$P(X_{n+1} = a | X_n = b) = P(X_{n+1} = a | X_n = b, X_1, \dots, X_{n-1})$$

"
 $P_{a,b}$

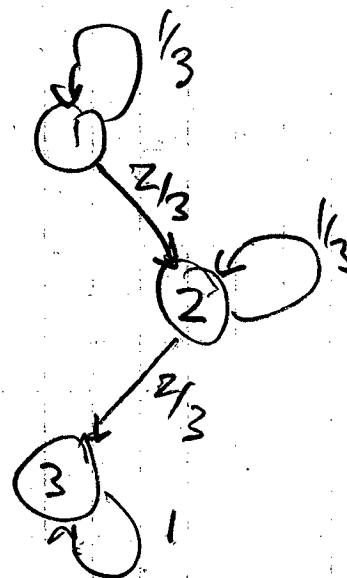
transition probability from a to b
(not a function of n)

Let's assume $X_i \in \{1, \dots, n\}$

$$\{P_{a,b}\}_{a=1, b=1}^{n,n}$$

transition matrix of the Markov chain

$$P = \begin{bmatrix} 1/3 & 2/3 & 0 \\ 0 & 1/3 & 2/3 \\ 0 & 0 & 1 \end{bmatrix}$$



$\rightarrow (1, 0, 0)$
 If $X_1 = 1, X_2 = \begin{cases} 2 & p = 2/3 \\ 1 & p = 1/3 \end{cases}$

If X_n has distribution v , then X_{n+1} has distribution vP

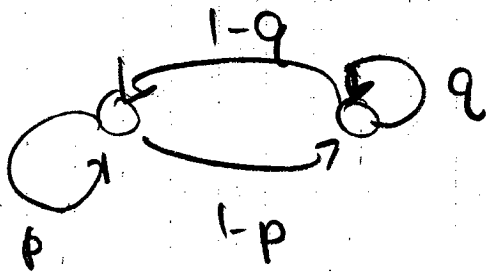
DEF

The stationary probability of a chain with transition matrix P is a probability distribution π such that

$$\pi P = \pi$$

For the example $\pi = (0, 0, 1)$

$$(0, 0, 1) \begin{pmatrix} 1/3 & 2/3 & 0 \\ 0 & 1/3 & 2/3 \\ 0 & 0 & 1 \end{pmatrix} = (0, 0, 1)$$



$$(x, 1-x) \begin{pmatrix} p & 1-p \\ 1-q & q \end{pmatrix} = (x, 1-x)$$

$$xp + (1-x)(1-q) = x$$

So, VP gives us the next state.

$$P \begin{pmatrix} 1 \\ \vdots \\ 1 \end{pmatrix} = \begin{pmatrix} 1 \\ \vdots \\ 1 \end{pmatrix} v$$

The all ones vector, $\vec{1}$, is a right-eigenvector of P (with eigenvalue 1)

In $\pi P = \pi$, π is a left-eigenvector (corresponding to eigenvalue 1).

Frobenius - Perron

Let P be a matrix with non-negative entries. Then P has a non-negative eigenvalue, λ , with non-negative eigenvector v .

For all other eigenvalues μ of P ,

$$|\mu| \leq \lambda.$$

If we replace "non-negative" with "positive" everywhere in the above, we get

$$|\mu| < \lambda.$$

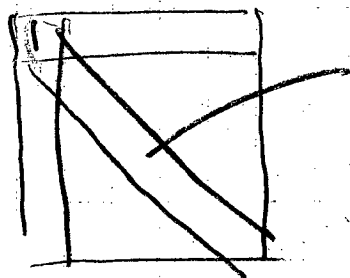
If we have a matrix P and we are computing P^m and can write $P = SDS^{-1}$, D a diagonal matrix, then

$$P^m = (SDS^{-1})^m = SD^m S^{-1}$$

Note that D is a diagonal matrix whose entries are the eigenvalues of P .

If we can't diagonalize P , then we can use Jordan Normal Form. We will not talk about this further.

If P has all entries positive (and we assume it's diagonalizable), how does D^n look?

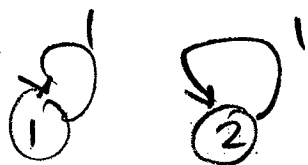


All of these are < 1 in absolute value. Therefore, the n^{th} powers go to 0 as $n \rightarrow \infty$.

This speaks to the stationary distribution.

The conclusion is that all entries in P positive implies there is a unique stationary distribution.

The follow is not unique:



$$x_1 = 1, x_2 = 1, \dots \quad \text{or} \quad x_1 = 2, x_2 = 2, \dots$$

$$\begin{bmatrix} x & y \end{bmatrix} \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} = \begin{bmatrix} x & y \end{bmatrix}$$

Note that not all entries in P are positive.

$$x + y = 1, \quad x, y \geq 0$$

It is sufficient to have 2 such that all entries in P^2 are positive.

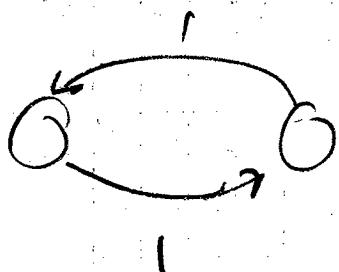
How do we achieve this? What are the obstacles

To have all entries in P^2 , we need

1. connected (for every pair of states a, b

$$\exists c_1, \dots, c_q \quad \left(\begin{array}{c} a \\ b \end{array} \right) \quad \left(\begin{array}{c} c_1 \\ c_2 \\ \vdots \\ c_q \end{array} \right) \quad \left(\begin{array}{c} p_{c_1, c_2} > 0 \\ p_{c_2, c_3} > 0 \\ \vdots \\ p_{c_{q-1}, c_q} > 0 \end{array} \right)$$

What about



0	1
1	0

$$(\frac{1}{2}, \frac{1}{2}) \quad P = (\frac{1}{2}, \frac{1}{2})$$

All entries in P positive

↓

There is a unique stationary distribution π and

$$\lim_{n \rightarrow \infty} \|V P^n - \pi\|_1 = 0 \quad \text{for any } V$$

So, we need

2. aperiodic

We avoid periodicity if $p_{v,v} > 0$ for all v .

Ergodic means connected and aperiodic.

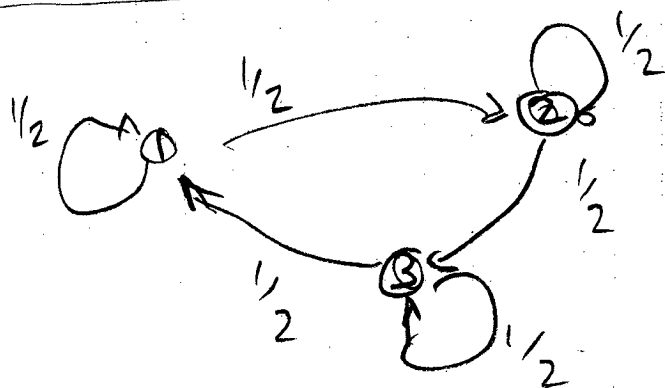
DEF

A Markov chain is reversible (let P be the transition matrix of an ergodic MC) if

$$\pi_i P_{ij} = \pi_j P_{ji} \quad (*)$$

↑ detailed balance condition

If P is ergodic and π satisfies the detailed balance condition, then π is a stationary distribution.



aperiodic $\left\{ \begin{array}{l} \Rightarrow \text{ergodic} \Rightarrow \text{unique stationary distribution} \end{array} \right.$

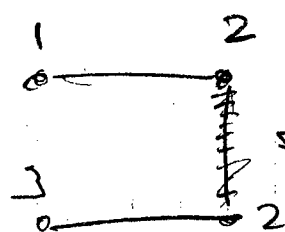
$$\pi = (1/3, 1/3, 1/3)$$

$$\pi_1 = \pi_2 = 1/3$$

$$P_{12} = 1/2, P_{21} = 0$$

$\therefore$ not reversible

Sample labelings of vertices of a graph
by 3 colors



$$P(\text{labeling}) \propto \lambda^{\# \text{ monochromatic edges}}$$

$$= \frac{1}{2} \lambda^{\# \text{ mono} \dots}$$

monochromatic

MC

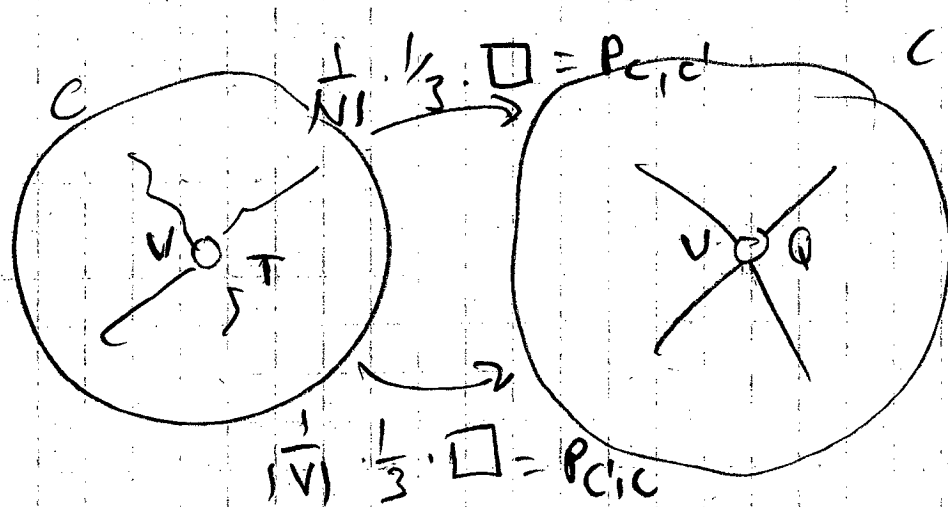
Current state $X_n = C$

Pick a random vertex v of G . Pick a random color $c' \in \{1, 2, 3\}$.

Let c' be coloring C with color of v changed to c' .

With probability $\square$ let $X_{n+1} = C'$, otherwise let $X_{n+1} = C$

Transitions in MC are differences in coloring of a single vertex, we want to satisfy the detailed balance condition



Goal $\pi_C P_{C,C'} = \pi_{C'} P_{C',C}$

The goal is

$$\frac{\pi_c}{\pi_{c'}} = \frac{p_{c'|c}}{p_{c|c'}}$$

$$\lambda = \frac{\# \text{ mono in } c - \# \text{ mono in } c'}{\lambda} = \frac{\square}{\square}$$

$$1 \leq \lambda$$

If c' has less monochromatic edges, then $\square = 1$, otherwise

$$\square = \lambda = \frac{\# \text{ mono in } c - \# \text{ mono in } c'}{\lambda}$$

Read about Metropolis's filter in the book. Don't need to read about extended version of Metropolis's.

$$\square = \begin{cases} 1 & \text{if } \pi_{c'} > \pi_c \\ \frac{\pi_{c'}}{\pi_c} & \text{otherwise} \end{cases}$$

Midterm Review

ML for multinomial distribution

$$p_k = p(x=k), \quad k=1, \dots, K$$

$$= \prod_{n=1}^N p(x_n | \theta) = \prod_{n=1}^N \prod_{k=1}^K p_k^{I\{x_n=k\}} \quad (\text{Bishop p 75})$$

Indicator function is 1 if $x_n = k$, 0 otherwise

Task is to find p_k that maximizes probability given observations x_n

$$\ln L = \sum_{n=1}^N \sum_{k=1}^K I\{x_n=k\} \ln p_k - \lambda \left(\sum_{k=1}^K p_k - 1 \right)$$

Remember the constraint that $\sum p_k = 1$

$$\frac{\partial \ln L}{\partial p_k} = \sum_{n=1}^K \frac{I\{x_n=k\}}{p_k} - \lambda = 0$$

$$\sum_{n=1}^N I\{x_n=k\} - \lambda p_k$$

Sum both sides over all k yields

$$p_k = \frac{\sum_{n=1}^N I\{x_n=k\}}{N} = \frac{N_k}{N}$$

Steps

1. Create an expression for seeing $x_1, \dots, x_n$
2. Eliminate π by taking log
3. Add constraints
4. Optimize by setting $\frac{\partial}{\partial \theta}$ to 0

Markov Chain

$$p^T = (p_1, \dots, p_N)$$

$$A \in \mathbb{R}^{N \times N}$$

$$p^T(t+1) = p^T(t) \cdot A$$

Solve for $p^T = p^T A$ to find the stationary distribution.

EM Algorithm

z hidden

x observed

θ parameters

we want to find $\underset{\theta}{\operatorname{argmax}} \phi(x|\theta)$

If this is too difficult, we look at $\phi(z, x|\theta)$.

E: $p(z|x, \theta^{\text{old}})$ $\leftarrow$ this is the form of q

M: $\theta^* = \underset{\theta}{\operatorname{argmax}} \sum p(z|x, \theta^{\text{old}}) \ln p(z, x|\theta)$

$$L(q, \theta) = \sum_z q(z) \ln \frac{p(x, z|\theta)}{q(z)}$$

Bernoulli EM

$$p(x) = \sum_{k=1}^K \pi_k p(x|\mu_k) = \prod_{k=1}^K [\pi_k p(x|\mu_k)]^{z_k}$$

$z_k = 1$ iff x is from the k^{th} mixture, otherwise 0

$$p(x|\mu_k) = \mu_k^x (1-\mu_k)^{1-x}$$

$$\textcircled{E} \gamma(z_{nk}) = \frac{\pi_k p(x_n|\mu_k)}{\sum_j \pi_j p(x_n|\mu_j)}$$

$\uparrow$
 n^{th} sample came
 from k^{th} class

$\textcircled{M.1}$

$$E_z[\ln p(z, x|\theta)] = E_z\left[\ln \prod_{n=1}^N \prod_{k=1}^K (\pi_k \cdot \mu_k^{x_n} \cdot (1-\mu_k)^{1-x_n})^{z_{nk}}\right]$$

$$= \sum_n \sum_k z_{nk} (\ln \pi_k + x_n \ln \mu_k + (1-x_n) \ln (1-\mu_k) - \lambda (\sum \pi_k - 1))$$

$$\frac{\partial E_z}{\partial \pi_k} = 0$$

$$\sum_n \frac{\partial (z_{nk})}{\pi_k} = \lambda$$

$$\sum_k \sum_n \partial(z_{nk}) = \lambda = N$$

$$\pi_k = \frac{N_k}{N}$$

For μ_k , we have

$$\frac{\partial E_z}{\partial \mu_k} = 0$$

we end up with $\sum_n \frac{\partial(z_{nk}) x_n}{\mu_k} = \sum_n \frac{\partial(z_{nk}) (1-x_n)}{1-\mu_k}$

$$\mu_k = \frac{\sum_n \partial(z_{nk}) x_n}{\sum_n \partial(z_{nk})}$$

4.2.2009

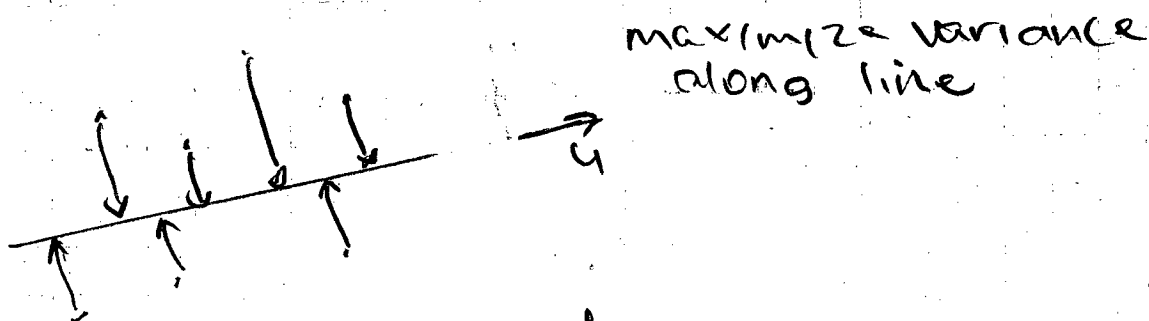
PCA

We have data points $x_1, \dots, x_n \in \mathbb{R}^d$ (d is "big")

We want to find points $y_1, \dots, y_n \in \mathbb{R}^M$ (M is "small")
capturing as much variance of x_i as possible

No labels, unsupervised

Case $M=1$



We want to find $u \in \mathbb{R}^d$, $\|u\|_2 = 1$

$$y_i = u^T x_i$$

Maximize variance of $u^T x_1, \dots, u^T x_n$, which is

$$\frac{1}{n} \sum_{i=1}^n (u^T x_i - u^T \bar{x})^2, \quad \bar{x} = \frac{1}{n} \sum_{i=1}^n x_i$$

$$V = \frac{1}{n} \sum_{i=1}^n u^T (x_i - \bar{x})(x_i - \bar{x})^T u$$

$$= u^T S u, \quad S = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})(x_i - \bar{x})^T \in \mathbb{R}^{d \times d}$$

We want to maximize $u^T S u$ subject to $\|u\|_2 = 1$.
We use Lagrange:

$$L(\lambda, u) = u^T S u - \lambda (u^T u - 1)$$

$$\frac{\partial L}{\partial u} = 2Su - 2\lambda u = 0$$

$$\boxed{Su = \lambda u}$$

u is an eigenvector of S

$$\text{The value of } u^T S u = u^T \lambda u = \lambda$$

The optimal choice of u is the eigenvector of S corresponding to the largest eigenvalue.

So for case $m=1$, we take $x_1, \dots, x_n$ and compute S

$$Su = \lambda u$$

$$\underline{m=2}$$

$$u_1 : Su_1 = \lambda_1 u_1$$

$$u_2 : Su_2 = \lambda_2 u_2$$

The eigenvalues of S satisfy $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n$

SVD

Let A be an $m \times n$ matrix. Then A can be written

$$A = U \Sigma V^T$$

where $U \in \mathbb{R}^{m \times m}$, $U^T U = I$
 $V \in \mathbb{R}^{n \times n}$, $V^T V = I$ (U, V are orthogonal)

$\Sigma \in \mathbb{R}^{m \times n}$ diagonal matrix with ≥ 0 entries.

The entries in Σ are called singular values of A .

$$\Sigma = \begin{bmatrix} \sigma_1 & & & 0 \\ & \sigma_2 & & \\ & & \ddots & \\ 0 & & & \sigma_m \\ & & & & 0 \end{bmatrix}, \quad \sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_m$$

$$\Sigma_T = \begin{bmatrix} \sigma_1 & & & 0 \\ & \sigma_t & & \\ & & \ddots & \\ & & & 0 \end{bmatrix}, \quad t < m$$

PCA

$$x_1, \dots, x_n \in \mathbb{R}^d$$

Project to $\mathbb{R}^m$ while maximizing variance of the projection.

Facts About SVD

We have A and want to approximate it with a matrix of rank k

$$\min_{B, \text{rk}(B) \leq k} \|A - B\|_2$$

$$\min_{B, \text{rk}(B) \leq k} \|A - B\|_F$$

$$\|A\|_2 := \max_{\|x\|_2=1} \|Ax\|_2 = (\rho(A^T A))^{1/2} = \sigma_1(A)$$

↖ largest eigenvalue

$$\|A\|_F := \sqrt{\sum_{i,j} A_{i,j}^2}$$

$$\boxed{A = U \Sigma V}$$

$$A - B = A - U \Sigma_k V$$

For $\text{rk}(A) = a$ and $\text{rk}(B) = b$, then
 $\text{rk}(AB) \leq \min\{a, b\}$

$$X = \begin{bmatrix} x_1^T - \bar{x}^T \\ x_2^T - \bar{x}^T \\ \vdots \\ x_n^T - \bar{x}^T \end{bmatrix}$$

this is a row ^{vector}, making
 X a matrix $\in \mathbb{R}^{n \times d}$

↓
SVD of X

$$X = U \Sigma V$$

$$X' = U \Sigma_M V$$

$$S = \frac{1}{2} X^T X$$

$M=1$ $\rightarrow$ largest eigenvalue of S and the corresponding eigenvector is the first column of V , $V \in \mathbb{R}^{d \times d}$

$$S = \frac{1}{n} (X^T X) = \frac{1}{n} (V^T \Sigma^T U^T U \Sigma V)$$

$$= \frac{1}{n} (V^T \Sigma^T \Sigma V)$$

$$\Sigma^T \Sigma = \begin{bmatrix} \sigma_1^2 & & \\ & \ddots & \\ & & \sigma_k^2 \end{bmatrix}$$

$$X = U \Sigma V$$

Use the first M columns of V for the projection.

$$\boxed{A} \cdot \boxed{B}$$

$$(a_i b_i)^T_{u,v} =$$

$$(a_i)_u (b_i)_v$$

$$\sum_{i=1}^n (a_i b_i^T)_{u,v}$$

$$= \sum_{i=1}^n \underbrace{(a_i)_u}_{A_{u,i}} \underbrace{(b_i)_v}_{B_{i,v}}$$

$$\begin{bmatrix} a_1 & \dots & a_n \\ b_1^T \\ \vdots \\ b_n^T \end{bmatrix}$$

$$X = \begin{pmatrix} x_1 - \bar{x}^T \\ \vdots \\ x_n - \bar{x}^T \end{pmatrix}$$

$n \times d$ matrix

goal SVD of X

eigenvalues
eigenvectors

of XX^T $n \times n$

eigenvalues
eigenvectors

of $X^T X$ $d \times d$

These two matrices have the same set of eigenvalues.

4.7.2009

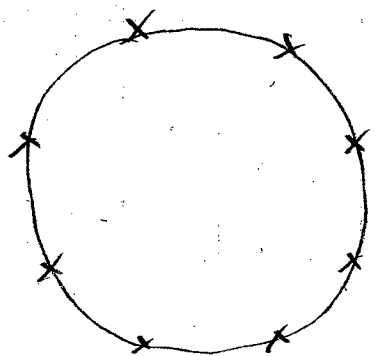
Manifolds

NON-LINEAR

Goal: Find a low dimensional structure in data.

(e.g. PCA was extracting a ^{LINEAR} subspace of dim M which contained most "information")

"We want to figure out low dimensional structures which are not necessarily linear."



data is lying on a circle

Imagine data lying on a bowl in 3-space or on the surface of a rolled-up piece of paper.

Assumption: The data comes from a low-dimensional manifold (embedded in a high-dimensional space)

Isomap(1) Look at data points $x_1, \dots, x_n \in \mathbb{R}^d$

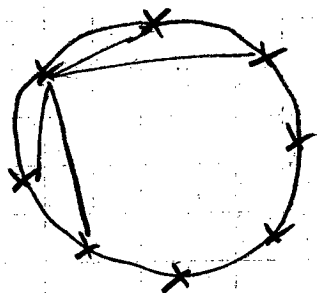
Create a graph. Vertices correspond to data points

Vertices i and j will be connected if

$$\|x_i - x_j\|_2 \leq \epsilon$$

Alternatively, connect i to $j_1, \dots, j_k$ where $x_{j_1}, \dots, x_{j_k}$ are k -NN of x_i

e.g.



We have a graph. Let's add edge weights

$$w_{ij} = \|x_i - x_j\|_2$$

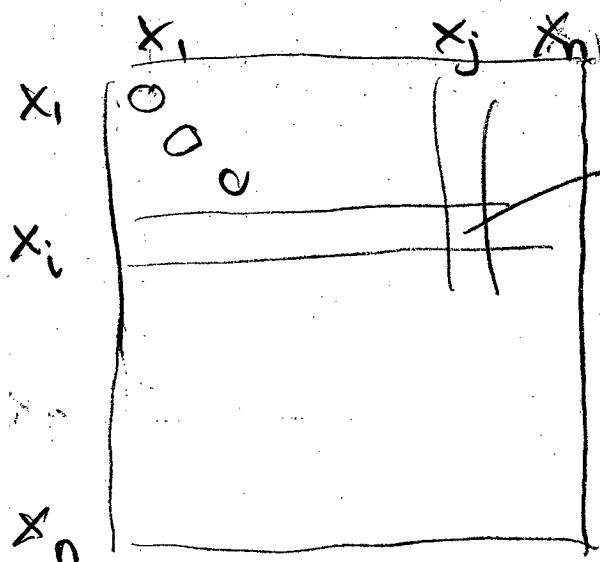
The algorithm works by finding geodesics on the manifold.

(2) "Recover" approximate geodesics"

(shortest curve between points)

Compute shortest path between every pair of vertices.

(3) Use the distances to embed the data into the low-dimensional Euclidean space



distance from i to j
in the graph
(symmetric)

We more or less preserve the distances.

MDS (Multidimensional Scaling)

Given distance matrix, find a good embedding in $\mathbb{R}^M$

Given distance matrix

↓
Find matrix of inner products

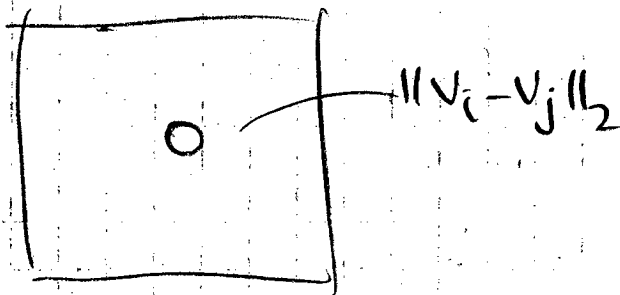
↓
SVD, Take top M eigenvectors

↓
Find good embedding in $\mathbb{R}^M$

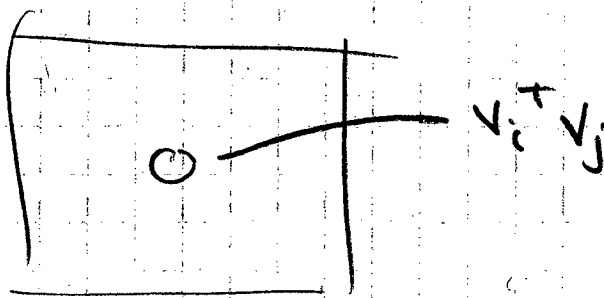
There are unknown $V_1, \dots, V_n$ in $\mathbb{R}^n$

$$\text{dist}(V_i, V_j) = \|V_i - V_j\|_2$$

We are given



We want



$$\begin{array}{c} A \\ \begin{bmatrix} 0 & \sqrt{2} & \sqrt{2} \\ \sqrt{2} & 0 & \sqrt{2} \\ \sqrt{2} & \sqrt{2} & 0 \end{bmatrix} \end{array} \longrightarrow \begin{array}{c} C \\ \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} \end{array}$$

Given matrix of distances, we want to calculate matrix of inner products assuming

$$\sum_{i=1}^n v_i = 0$$

i.e. The points are centered around a mean.

A. Given $\|v_i - v_j\|_2$ ($\forall i, j$)

B. Find $v_i^T v_j$ ($\forall i, j$)

$$B_{ij} = A_{ij}^2 = \|v_i - v_j\|_2^2$$

Given C , $B_{ij} = C_{ii} + C_{jj} - 2C_{ij} = (v_i - v_j)^T (v_i - v_j)$

How do we compute C from A ?

$$C = -\frac{1}{2} H B H \quad \text{where} \quad H = I - \frac{1}{n} J, \quad J = \mathbf{1} \mathbf{1}^T$$

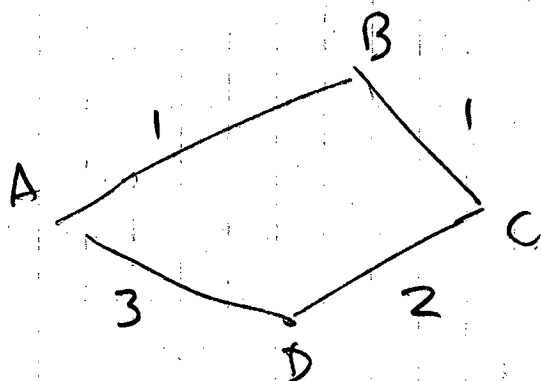
$$\begin{array}{l} J = \begin{bmatrix} 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \end{bmatrix} \\ J = \begin{bmatrix} 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \end{bmatrix} \end{array}$$

Laplacian Eigenmaps

DEF

Given a graph $G = (V, E)$ with weighted edges
 (Let w_{ij} = weight of edge $\{i, j\}$, 0 if not an edge)

NON-NEGATIVE



$$\begin{array}{c|cccc} & A & B & C & D \\ \hline A & 0 & 1 & 0 & 3 \\ B & 1 & 0 & 1 & 0 \\ C & 0 & 1 & 0 & 2 \\ D & 3 & 0 & 2 & 0 \end{array}$$

 $= W$

$$\begin{pmatrix} 4 & 0 & 0 & 0 \\ 0 & 2 & 0 & 0 \\ 0 & 0 & 3 & 0 \\ 0 & 0 & 0 & 5 \end{pmatrix} = D$$

Let $D = \text{diagonal} \left(\sum_{j=1}^n w_{ij} \right)$

Laplacian matrix of G is $L = D - W$

$$\begin{pmatrix} 4 & -1 & 0 & -3 \\ -1 & 2 & -1 & 0 \\ 0 & -1 & 3 & -2 \\ -3 & 0 & -2 & 5 \end{pmatrix} = L$$

Fact

$$0 \leq \sum_{\{i,j\} \in E} w_{ij} (y_i - y_j)^2 = \mathbf{z}^T \mathbf{L} \mathbf{y} \Rightarrow L \text{ is positive semidefinite}$$

$$\begin{aligned}\sum_{(i,j)} w_{ij} (y_i - y_j)^2 &= \sum_{(i,j)} w_{ij} y_i^2 + \sum_{(i,j)} w_{ij} y_j^2 - 2 \sum_{(i,j)} w_{ij} y_i y_j \\ &= \sum_i D_i y_i^2 + \sum_j D_j y_j^2 - 2 \sum_{(i,j)} w_{ij} y_i y_j\end{aligned}$$

$$\begin{aligned}2Y^T L Y &= 2 \sum_{(i,j)} y_i y_j L_{ij} \\ &= -2 \sum_{(i,j)} y_i y_j w_{ij} + 2 \sum_{i=1}^n y_i^2 D_i\end{aligned}$$

take data points $x_1, \dots, x_n$

(1) Create graph (either k-NN or $\leq \epsilon$ distance)

$$w_{ij} = \begin{cases} 1 & \text{if } i, j \text{ is connected} \\ \emptyset & \text{otherwise} \end{cases}$$

$$w_{ij} = \begin{cases} \frac{1}{\|x_i - x_j\|_2^2} & \text{if } i, j \text{ is connected} \\ \emptyset & \text{otherwise} \end{cases}$$

(2) Compute L , the Laplacian of G

(3) This will be $\emptyset$
 ~~$L V_1 = \lambda_1 D V_1$~~

$$0 = \lambda_1 \leq \lambda_2 \leq \dots \leq \lambda_{M+1}$$

$M+1$ smallest eigenval.
of L

$$L V_{M+1} = \lambda_{M+1} D V_{M+1}$$

$$x \mapsto (x^T v_2, \dots, x^T v_{M+1})$$

$$0 \leq \sum_{\{i,j\} \in E} w_{ij} (y_i - y_j)^2 = 2y^T Ly$$

We want to enforce $y^T \mathbf{1} = 0$

$$y^T D y = 1$$

Hidden Markov Models (HMMs)

Markov chain (time homogeneous)

Sequence of random variables $X_1, \dots, X_n, \dots$

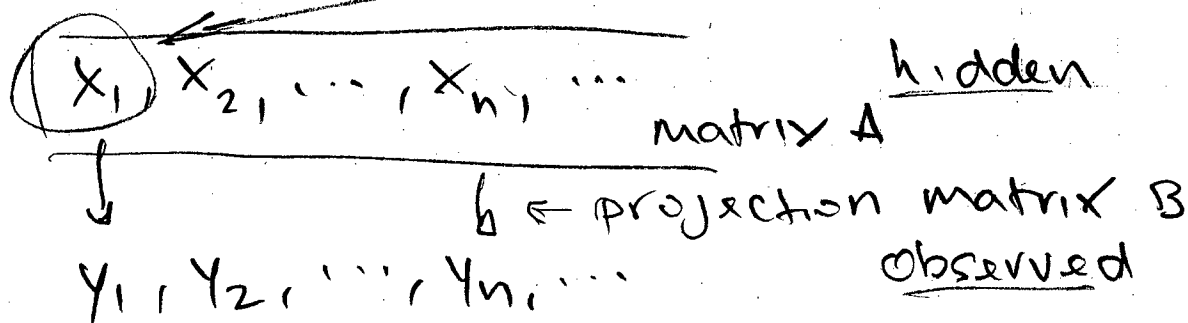
$$P(X_{n+1}=a \mid X_n=b, X_{n-1}, \dots, X_1) = P(X_{n+1}=a \mid X_n=b) = A_{ab}$$

Time homogeneous means that it is not dependent on n

$$X_i \in \Omega$$

$$A \in \mathbb{R}^{|\Omega| \times |\Omega|}$$

Need to know π = distribution of X_1



$$P(Y_i=c \mid X_i=d) = B_{c,d} = P(Y_i=c \mid X_i=d, X_1, \dots, X_{i-1})$$

An HMM is given by (M hidden states, N observable states)

π distribution on $[M]$

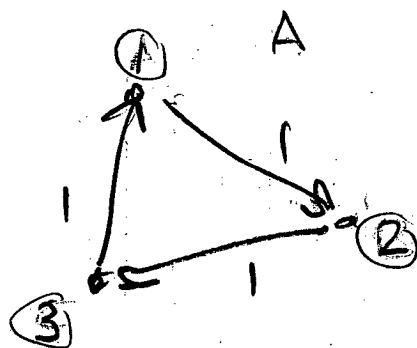
$$A \in \mathbb{R}^{M \times M}$$

$$B \in \mathbb{R}^{M \times N}$$

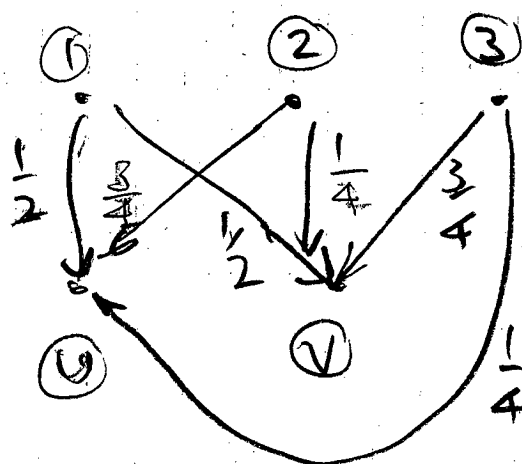
Ex.

$$\pi = (\frac{1}{3}, \frac{1}{3}, \frac{1}{3})$$

116



B.



hidden

visible

you see a sequence of U and V , you want to discover the sequence of hidden states

DECODING (ML)

1) Given A, B, π , and $y_1, \dots, y_T$

find $x_1, \dots, x_T$ which maximizes the likelihood.

DECODING (MAP)

2) Given $A, B, \pi, y_1, \dots, y_T$

Find posterior for x_i

EVALUATION3. Given $A, B, \pi, y_1, \dots, y_T$

What is the probability of this sequence?

LEARNING4. Given $y_1^{(c)}, \dots, y_T^{(c)}$

$$\vdots$$

$$y_1^{(k)}, \dots, y_T^{(k)}$$
Find A, B, π which maximize the likelihood.Decoding

1) Viterbi's Algorithm

o dynamic programming

	1	...	T
1			
⋮			
M			

Table D

$$D_{ij} = \text{maximum probability of a sequence } x_1, \dots, x_j \text{ with } x_j = i$$

Take all sequences of internal states of length j which end with i . Take the one with the highest probability. D_{ij} is its probability.

$$\text{Prob}(x_1, \dots, x_T | y_1, \dots, y_T) = \pi(x_1) \left(\prod_{i=1}^{T-1} A_{x_i, x_{i+1}} \right) \left(\prod_{i=1}^T B_{x_i, y_i} \right)$$

We calculate D_{ij} by first determining $D_{k,j-1}$ for $k \in [M]$

$$D_{iN} = \max_{k \in [M]} D_{k,j-1} \cdot A_{k,i} \cdot B_{x_i, y_j}$$

The first column of the table is initialized with

$$\pi_i B_{i, y_1}$$

The running time is $(MT) M = M^2 T$

of entries

time per entry

Evaluation

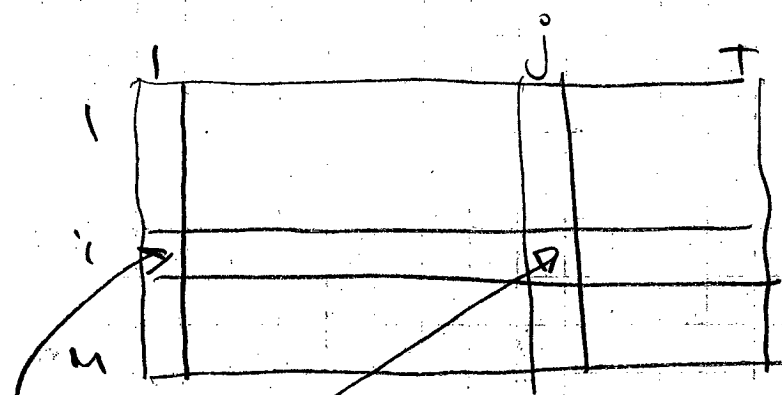
3) Dynamic programming

We want to compute

$$P(y_1, \dots, y_T | A, B, \pi) =$$

$$\sum_{x_1, \dots, x_T=1}^M P(y_1, \dots, y_T, x_1, \dots, x_T | A, B, \pi)$$

This is actually T sums.



$$D_{i,j} = P(Y_1, \dots, Y_j, X_j = i | A, B, \pi)$$

$$D_{i,1} = \pi_i B_{i,Y_1} =$$

The recurrence relation is

$$D_{i,j} = \sum_{k=1}^M D_{k,j-1} \cdot A_{k,i} \cdot B_{i,Y_j}$$

~~(Note that k is the value to which we condition X_{j-1} .)~~

$$P(Y_1, \dots, Y_j | X_j = i) \stackrel{?}{=} \sum_{k=1}^M P(Y_1, \dots, Y_{j-1} | X_j = k) A_{k,i} B_{i,Y_j}$$

~~↑ is this true~~

$$P(Y_1, \dots, Y_j | X_j = i) = \frac{P(Y_1, \dots, Y_j, X_j = i)}{P(X_j = i)}$$

$$\sum_{k=1}^M (\text{from above}) = \sum_{k=1}^M \frac{P(Y_1, \dots, Y_{j-1}, X_{j-1} = k)}{P(X_{j-1} = k)} \cdot A_{k,i} \cdot B_{i,Y_j}$$

We got confused, so we changed to

$$D_{i,j} = P(y_1, \dots, y_j, X_j = i | A, B, \pi)$$

$$D_{i,1} = \pi_i B_{i,y_1}$$

$$P(y_1, \dots, y_j, X_j = i) = \sum_{k=1}^M P(y_1, \dots, y_{j-1}, X_{j-1} = k) A_{k,i} B_{i,y_j}$$

The answer is the sum of the last (T) column:

$$P(y_1, \dots, y_T | A, B, \pi) = \sum_{i=1}^M D_{i,T}$$

2)

$$P(X_i = q | y_1, \dots, y_T, A, B, \pi) = ?$$

$$P(X_i = q | y_1, \dots, y_T, A, B, \pi)$$

$$P(y_{i+1}, \dots, y_T | X_i = q, A, B, \pi)$$

$$= \frac{P(y_1, \dots, y_T | X_i = q) P(X_i = q)}{P(y_1, \dots, y_T)}$$

leaving
A, B, π out
for
brevity

HMM

transition matrix $A \in \mathbb{R}^{M \times M}$
 initial probability $\pi \in \mathbb{R}^M$
 emission matrix $B \in \mathbb{R}^{M \times N}$

$$X_i \sim \pi \quad P(X_{i+1}=a | X_i=b, \dots, X_1) = A_{a,b}$$

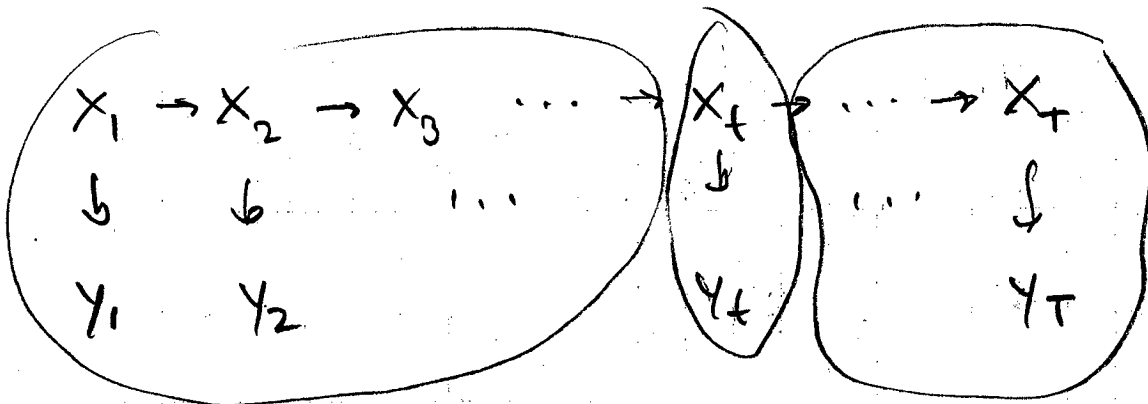
$$P(Y_i=c | X_i=d) = B_{c,d}$$

Q2 Given $Y_1, \dots, Y_T$, what is the likelihood of X_t

$$P(X_t | Y_1, \dots, Y_T)$$

We know A, B, π .

$$P(X_t | Y_1, \dots, Y_T) = \frac{P(Y_1, \dots, Y_T | X_t) P(X_t)}{P(Y_1, \dots, Y_T)} = *$$



$$P(Y_1, \dots, Y_T | X_t) = P(Y_1, \dots, Y_t | X_t) P(Y_{t+1}, \dots, Y_T | X_t)$$

$$* = \frac{P(Y_1, \dots, Y_t | X_t) P(Y_{t+1}, \dots, Y_T | X_t)}{P(Y_1, \dots, Y_T)} = \alpha_t(a) \beta_t(a)$$

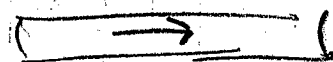
Let $X_t = a$

$$P(y_1, \dots, y_T) = \sum_{a \in [M]} \alpha_t(a) \beta_t(a)$$

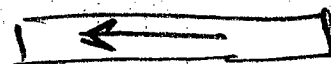
$\alpha_t(a)$ = probability of seeing $y_1, \dots, y_t$ and ending up in $X_t = a$

$\beta_t(a)$ = probability of seeing $y_{t+1}, \dots, y_T$ given that $X_t = a$

$$\alpha_{t+1}(a') = \sum_{a \in [M]} \alpha_t(a) A_{a,a'} \beta_{a', y_{t+1}}$$



$$\beta_t(a) = \sum_{a' \in [M]} \beta_{t+1}(a') A_{a,a'} \beta_{a, y_t}$$



$$\alpha_1(a) = \pi_a \beta_{a, y_1}$$

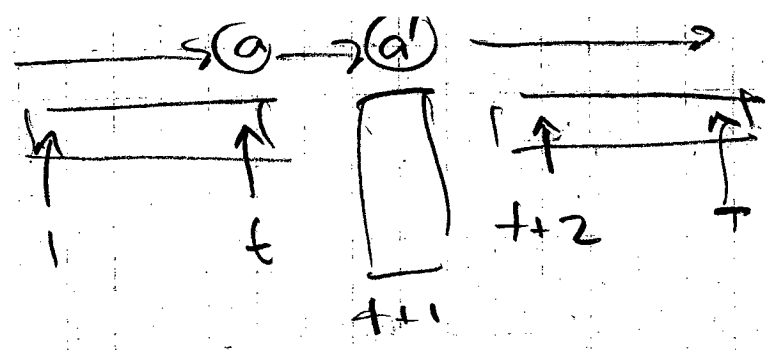
$$\beta_{T+1}(a) = \beta_{a, y_T} \quad \beta_T(a) = 1$$

Use recurrence, we get (as a check)

$$\beta_{T+1}(a) = \sum_{a' \in [M]} \underbrace{\beta_T(a')}_{1} A_{a,a'} \beta_{a, y_T} = \beta_{a, y_T}$$

$$P(X_t = a, X_{t+1} = a' | y_1, \dots, y_T) = \alpha_t(a) \beta_{t+1}(a')$$

$$\alpha_t(a) \beta_{t+1}(a') A_{a,a'} \beta_{a', y_{t+1}}$$



Q4 Learning

We get $y_1^{(1)}, \dots, y_T^{(1)}$

$y_1^{(k)}, \dots, y_T^{(k)}$

We want to learn A, B, π .

We are going to use the EM algorithm.

Initialize A, B, π ,
Update A, B, π using "inferred internal state"

Compute

$$P(x_t = a, x_{t+1} = a' | y_1^{(k)}, \dots, y_T^{(k)}, A, B, \pi), k=1, \dots, K$$

$$\sum_{k=1}^K \sum_{t=1}^{T-1} P(\dots) = A'_{a,a'}$$

Then normalize rows of $A'_{a,a'}$ so that they sum to one.

$$\sum_{k=1}^K \sum_{t: y_t=b} P(x_t = a | y_1^{(k)}, \dots, y_T^{(k)}, A, B, \pi) = B'_{a,b}$$

Then normalize rows of $B'_{a,b}$ to sum to one
; means "such that"

$$\pi'(a) \propto \sum_{\ell=1}^K P(X_1 = a | Y_1^{(\ell)}, \dots, Y_T^{(\ell)} | A, B, \pi)$$

What does the likelihood function for this EM problem look like?

$$P(X_1, \dots, X_T, Y_1, \dots, Y_T) =$$

$$\pi_{X_1} \prod_{t=1}^{T-1} A_{X_t, X_{t+1}} \prod_{t=1}^T B_{X_t, Y_t}$$

Markov Decision Process (MDP) = (S, A, P_a, R, \gamma)

S = Set of states

A = actions

$P_a(S, S')$ = probability of going from S to S' (with action A)

$R: S \rightarrow R$ = reward

γ = discount factor (inflation), $\gamma < 1$

Given S, A, P_a, R, γ , the goal is to figure out what to do.

Determine a policy $\pi: S \rightarrow A$
 good

Initial state is $z \in S$

$$X_1 = z, X_2, \dots, X_t, \dots$$

$$P(X_{t+1} = S' | X_t = S) = P_{\pi_S}(S, S')$$

Reward in instantiation $x_1, \dots, x_t, \dots$ is
 $R(x_1) + \gamma R(x_2) + \gamma^2 R(x_3) + \dots = \sum_{t=0}^{\infty} \gamma^t R(x_{t+1})$

We want to maximize

$$E \left[\sum_{t=0}^{\infty} \gamma^t R(x_{t+1}) \right] \text{ by picking } \pi$$

4.15.2009

From HW # 4.3

If x and y are random variables, and $z = x + y$ then the distribution of z is the convolution of the distribution of x by the distribution of y .

$$p(z=a) = \sum_k p(x=k) p(y=a-k)$$

this is a convolution

VC-dimension of a circle is 3. Prove by setting 2 most distant points to '+', others to '-'.

Prove that $\exists d \pi_c(m) = O(m^d) \Rightarrow VC(C) < \infty$.

If $VC(C) = \infty$ then $\forall m \pi_c(m) = 2^m \Rightarrow \nexists d \forall m 2^m = O(m^d)$.
 (Proof by contradiction)

1. Prove that, given a set of n i.i.d. observations $x_1, \dots, x_n$ a distribution q that minimizes the KL-divergence with empirical distribution p , also maximizes the probability of the data $q(x_1, \dots, x_n)$
2. Derive (w/o using VC-dim) a PAC-learning bound for the concept class of axis-aligned SQUARES.

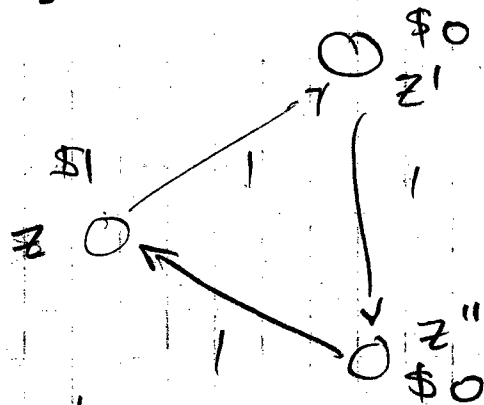
4.16.2009

Markov Rewards Process $S, \mathcal{A}, R, P, \gamma$ no
actionsStart in state s , make transition according to P , repeat, in each state get discounted reward. $x_0, x_1, \dots, x_n, \dots$ $x_0 = z$

$$P(x_{i+1} = a \mid x_i = b, \dots, x_0) = P_{a,b}$$

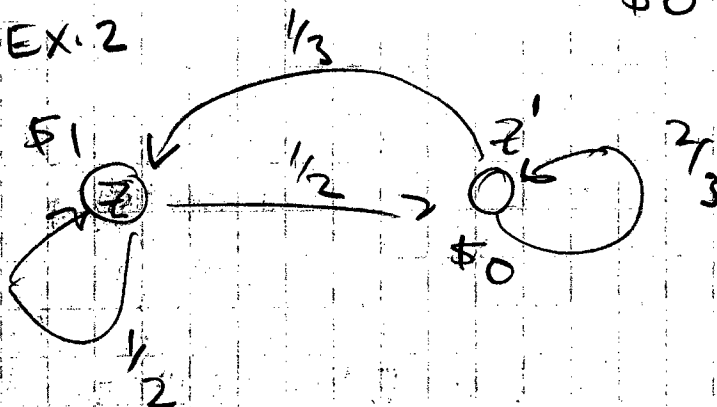
$$E\left[\sum_{i=0}^{\infty} \gamma^i R(x_i)\right] = \text{expected total reward}$$

EX. 1

 $\gamma = 0.99$

$$\frac{1}{1 - (0.99)^3}$$

EX. 2

 $V_s = \text{expected reward at } s$

$$(V_s) = \left(\gamma \sum_{s' \neq s} P(s, s') V_{s'} \right) + R(s)$$

For previous Ex. 1:

$$\begin{cases} V_z = 1 + \gamma V_{z'} \\ V_{z'} = \gamma V_{z''} \\ V_{z''} = \gamma V_z \end{cases} \quad V_z = 1 + \gamma^3 V_z$$

For Ex. 2:

$$V_z = 1 + \left(\frac{1}{2} V_z + \frac{1}{2} V_{z'} \right) \gamma$$

$$V_{z'} = \left(\frac{2}{3} V_{z'} + \frac{1}{3} V_z \right) \gamma$$

$(V_1, \dots, V_n)$ = vector of expected utilities = V^T

The matrix form for the equation for V_s is

$$V^T = \gamma V^T P + R^T$$

where $R^T = (R(1), \dots, R(n))$.

$$V^T (1 - \gamma P) = R^T$$

$$V^T = \underbrace{(1 - \gamma P)^{-1}}_{\text{is this always invertible}} R^T$$

is this always invertible

P satisfies $P_{ij} \geq 0$, $\sum_{j=1}^n P_{ij} = 1$. P is a stochastic matrix

$$\rho(P) = 1$$

(absolute value of)
spectral radius = largest λ eigenvalue of P

Largest eigenvalue of $\gamma P = \gamma < 1$

$$AV_i = \lambda_i V_i$$

$$(I - A)V_i = V_i - \lambda V_i = (1 - \lambda)V_i$$

A eigenvalues $\lambda_1, \dots, \lambda_n$

$I - A$ eigenvalues $1 - \lambda_1, \dots, 1 - \lambda_n$

Since $0 < \lambda_i < 1$, $0 < 1 - \lambda_i < 1$.
Therefore A is not singular

Markov Decision Process

$S, A, \{P_a, a \in A\}, R, \gamma$

We study policies (functions $S \rightarrow A$)

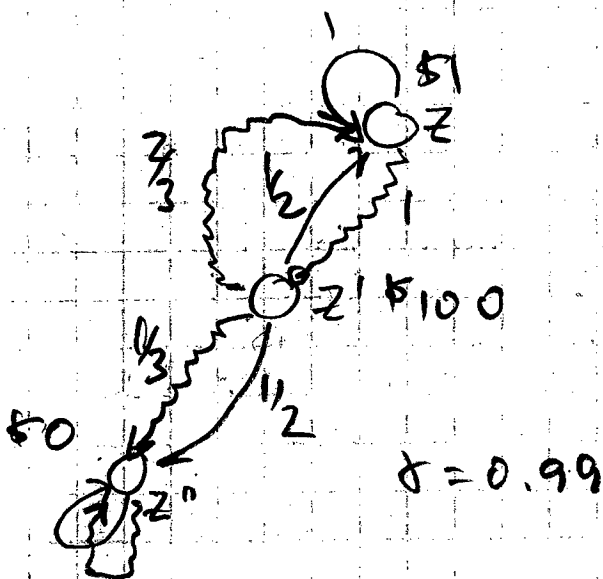
Given a policy π , the utility is

$$E \left[\sum_{i=0}^{\infty} \gamma^i R(x_i) \right]$$

where

$x_0, x_1, \dots, x_n, \dots$

$$P(x_{i+1} = s' | x_i = s, \dots) = P_{\pi(s)}(s, s')$$



policy a ———

policy b ~~~~~

There always exists an optimal policy.

How do we find optimal policies?

What is the utility of state s under optimal policy π^* ?

$$V_s = R(s) + \gamma \left(\max_a \sum_{s' \in S} P_a(s, s') V_{s'} \right)$$

(Bellman's equations)

~~Thus~~ Bellman's equations have a unique solution.

Value Iteration

Start with some V 's, e.g. $V^{(0)}(s) = 0$ ($\forall s \in S$)

Repeat

$$V^{(t+1)}(s) = R(s) + \gamma \max_a \sum_{s' \in S} P_a(s, s') V^{(t)}(s')$$

We hope this converges.

V_1

V_2

↓ value iteration

↓ value iteration

V_1'

V_2'

$$\|V_1 - V_2\|_\infty \leq \Delta$$

$1, 2$

$$\forall s \quad |V_1(s) - V_2(s)| \leq \Delta$$

$$\|V_1' - V_2'\|_\infty \leq ? \quad \Delta \gamma$$

$$V_1'(s) - V_2'(s) = \gamma \max_a \sum_{s' \in S} P_a(s, s') V_1(s') - \gamma \max_{a'} \sum_{s' \in S} P_{a'}(s, s') V_2(s')$$

$$\leq \gamma \max_a \sum_{s' \in S} P_a(s, s') (V_1(s') - V_2(s')) \leq \Delta \gamma$$

We get this
by choosing
suboptimal $a' = a$

Banach's Contraction Theorem

$$F: B \rightarrow B$$

$$\|f(x) - f(y)\| \leq \gamma \|x - y\| \Rightarrow \text{unique fixed point}$$

If a mapping is contracting, it has a unique fixed point.

V^* = fixpoint (utilities for optimal policy)

$$\|V(0) - V^*\|_\infty \leq ?$$

$$\|V(T) - V^*\|_\infty \leq 2\gamma^T$$

$$V_s = R(s) + \gamma \max_{a \in A} \sum_{s'} P_a(s, s') V_{s'}$$

Find the solution of the above equation.

$$V_s \geq R(s) + \gamma \sum_{s'} P_a(s, s') V_{s'} \quad (\text{for all } a \in A, s \in S)$$

How do we achieve equality? We add an objective to push V_s as low as possible

Linear Program

$$\min \sum_{s \in S} V_s$$

We claim that the optimal solution of the linear program is the unique solution to Bellman's equation

Policy Iteration

Start with policy $\pi^{(0)}$

Take the Markov Rewards Process for $\pi^{(t)}$
 Compute its utilities $v^{\pi^{(t)}}$
 (solving a system of equations)

Greedily update the policy.

$$\pi^{(t+1)}(s) = \operatorname{argmax}_a \sum_{s'} P_a(s, s') v^{\pi^{(t)}}(s')$$

Take action a from state s and then follow $\pi^{(t)}$

You could put $R(s)$ and γ in here, but it doesn't effect the result.

Thm 1) Policy improvement theorem

$$v^{\pi^{(t+1)}}(s) \geq v^{\pi^{(t)}}(s)$$

Thm 2) one shot improvement theorem

If $\pi^{(t)}$ is not optimal, then

$$\pi^{(t+1)} \neq \pi^{(t)}$$

4.21.2009

MDP

$$S, A, P(s'|s, a), R, \gamma$$

$$\text{policy } \pi: S \rightarrow A$$

$$X_0, X_1, \dots, X_t \quad X_0 = S$$

$$P(X_{t+1} = s' | X_t = s) = P(s' | s, \pi(s))$$

$$\text{We want to maximize } E \left[\sum_{t=0}^{\infty} \gamma^t R(X_t) \right]$$

Last time, we did value iteration:

$$V_0 \xrightarrow{T} V_1 \xrightarrow{T} \dots \xrightarrow{T} V^*$$

$$(TV)(s) = R(s) + \gamma \max_{a \in A} \sum_{s' \in S} P(s'|a, s) V(s')$$

T is a contraction.

$$\|TV_1 - TV_2\|_{\infty} \leq \gamma \|V_1 - V_2\|_{\infty}$$

take a policy π . Define $T_{\pi}: \mathbb{R}^S \rightarrow \mathbb{R}^S$

$$(T_{\pi}V)(s) = R(s) + \gamma \sum_{s' \in S} P(s'|s, \pi(s)) V(s')$$

Is T_{π} a contraction?

$$\|V_1 - V_2\|_{\infty} \leq \Delta$$

$$\begin{aligned} TV_1(s) - TV_2(s) &= \gamma \sum_{s'} P(s'|s, \pi(s)) (V_1(s') - V_2(s')) \\ &\leq \gamma \sum_{s'} P(s'|s, \pi(s)) \Delta = \Delta \gamma \end{aligned}$$

Yes, it is a contraction. What is the fixed point?

$$V_0 \xrightarrow{T_\pi} V_1 \xrightarrow{T_\pi} \dots \xrightarrow{T_\pi} V_\pi^\pi, \quad V_\pi(s) = E \left[\text{reward under policy } \pi \mid x_0 = s \right]$$

Markov Rewards Process is the same as MDP with a fixed policy. So, above is iterative solution to linear equations from MDP

DEF: π is optimal if $V^\pi = V^*$

Lemma: π is optimal if $T_\pi V^* = V^*$

Policy Iteration

$$\pi_0 \rightarrow \pi_1 \rightarrow \dots \rightarrow$$

↓ compute V^{π_0} (fixed point of T_{π_0})

$$\pi_1(s) = \operatorname{argmax}_a \sum_{s'} P(s'|s,a) V^{\pi_0}(s')$$

Lemma

$$V^{\pi_{t+1}} \geq V^{\pi_t} \quad (\text{i.e. } (\forall s \in S) \quad V^{\pi_{t+1}}(s) \geq V^{\pi_t}(s))$$

For value iteration, running time is dependent on δ . For policy iteration, running time is not dependent on δ .

$$V^{\pi_{t+1}}, \dots, \pi_{t+1}, \pi_t, \dots, \pi_t, \dots$$

$$\Downarrow$$

$$V^{k \times \pi_{t+1} + (10-k) \times \pi_t}$$

Apply π_{t+1} for first k steps, then π_t forever.

Proof by Induction

$$V^{k \times \pi_{t+1} + (10-k) \times \pi_t} \geq V^{\pi_t} \quad (t \text{ is fixed})$$

Base case: $k=0 \quad V^{\pi_t} \geq V^{\pi_t}$

$$\begin{aligned}
 V^{\pi_{t+1}}(s) &= \max_{a \in A} \sum_{s'} P(s'|a, s) V^{\pi_t}(s') \\
 &\geq \sum_{s'} P(s'|\pi_t(s), s) V^{\pi_t}(s')
 \end{aligned}$$

$$V^{(\infty)}(s) = \sum_{s'} P(s'|\pi_{t+1}(s), s) V^{\pi_{t+1}, \pi_t}(s') =$$

$$\sum_{s'} P(s'|\pi_{t+1}(s), s) V^{\pi_{t+1}, \pi_t}(s') \geq$$

$$\sum_{s'} P(s'|\pi_{t+1}(s), s) V^{\pi_t}(s') =$$

$$\begin{aligned}
 \max_a \sum_{s'} P(s'|a, s) V^{\pi_t}(s') &\geq \sum_{s'} P(s'|\pi_t(s), s) V^{\pi_t}(s') \\
 &= V^{\pi_t}(s)
 \end{aligned}$$

For induction, replace z by $k+1 \dots$

So, now, we know that

$$V_{\pi_0} \leq V_{\pi_1} \leq \dots \leq V_{\pi_t} \leq V_{\pi_{t+1}} \leq \dots$$

There are at most $|S|^{|A|}$ different policies (different valuation vectors)

There exists $t \leq |S|^{|A|}$, such that $V_{\pi_t} = V_{\pi_{t+1}}$

At this point, policy iteration stops and declares that π_t is the optimal policy.

How do we prove this?

$$\begin{aligned}
 T_{\pi_{t+1}} V_{\pi_t} &= T_{\pi_{t+1}} V_{\pi_t} \quad \text{— This is the definition of } T_{\pi_{t+1}} \\
 &\quad \left(\text{picks the optimal greedy action for any } t \right)
 \end{aligned}$$

$$\begin{array}{ccc}
 T_{\pi_{t+1}} & V_{\pi_t} & = TV_{\pi_t} \\
 & \text{"} & \text{"} \\
 T_{\pi_{t+1}} & V_{\pi_{t+1}} & \\
 & \text{"} & \text{"} \\
 & V_{\pi_{t+1}} & = TV_{\pi_{t+1}}
 \end{array}$$

Since $V_{\pi_{t+1}}$ is the fixed point of $T_{\pi_{t+1}}$

$$V_{\pi_{t+1}} = TV_{\pi_{t+1}} \Rightarrow V_{\pi_{t+1}} = V^*$$

What do you do if you don't know what the rewards are in advance?

We don't know $P(s|s', a)$, $R(s)$.

There is an underlying MDP with S, A, P, R, γ .
We know S, A, γ .

Q-Learning

$$Q(s, a) = R(s) + \sum_{s' \in S} P(s'|s, a) V(s')$$

If we knew Q , what is the optimal π ?

$$V(s) = \max_a Q(s, a)$$

$$\pi(s) = \operatorname{argmax}_a Q(s, a)$$

α_t decreases with time (as in simulated annealing)

Goal: Figure out Q .

$$\hat{Q}(s, a) \leftarrow (1 - \alpha_t) \hat{Q}(s, a) + \alpha_t (R(s) + \gamma \max_{a'} Q(s', a'))$$

In state s , we pick action a and end up in state s'

Thm. If $\sum \alpha_t = \infty$, $\sum \alpha_t^2 < \infty$, and every state visited infinitely often, then $\hat{Q} \rightarrow Q$

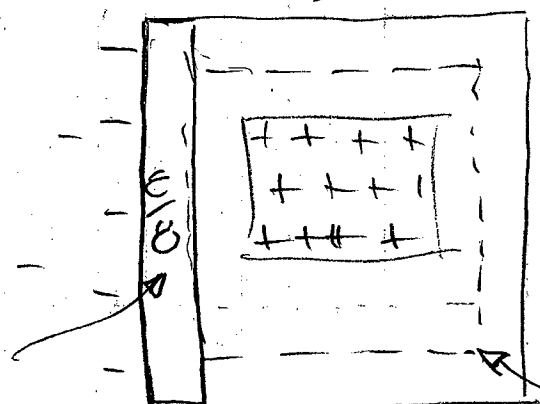
ϵ random action

$1-\epsilon$ action $a = \operatorname{argmax}_a q(s, a)$

4.22.2009

PAC learning for squares from 4.15

Must use negatives as well as positives



(stripes)
Eight regions
with $\epsilon/8$ error

or
this
stripe

--- actual square

KL divergence from 4.15

$$D(p \parallel q) = \sum p(x) \log \frac{p(x)}{q(x)} = \sum p(x) \log p(x) - \sum p(x) \log q(x)$$

$$p(x) = \frac{n_x}{N} \quad \sum_x n_x = N$$

$$= \frac{\sum_i \mathbb{I}\{x_i = x\}}{N}$$

$$D(p \parallel q) = \text{const} - \sum \frac{n_x}{N} \log q(x)$$

$$\prod_i p(x_i) = \prod_x q^{n_x}(x)$$

$$\begin{aligned} \operatorname{argmax}_{\phi} \prod_i \phi(x_i) &= \operatorname{argmax}_{\phi} \prod_x \phi^{n_x}(x) = \operatorname{argmax}_{\phi} \sum \ln \phi^{n_x}(x) \\ &= \operatorname{argmax}_{\phi} \sum n_x \ln \phi(x) \end{aligned}$$

EM Algorithm (from M. term problem)

E-step: "empirical assignment of data points to clusters"

$$P(x_i | z=1) \propto \frac{1}{10}$$

$$P(x_i | z=2) \propto \frac{1}{11} \quad \dots$$

$$P(x_i | z=3) \propto \frac{1}{12}$$

M-step: "setting α_i and β_i to maximize 'L'"

$$L = \sum_{x_i, z} P(z=z | x_i, \theta) \log P(z=z | x_i, \theta^*)$$

E-step

M-step: Find θ^* maximizing L

$$\log P(z, x | \theta^*) = \begin{cases} -10 & \text{if } \alpha_2 < x \\ \log \frac{\beta_2}{\alpha_2} & \text{if } \alpha_2 \geq x \end{cases}$$

$$L = \sum_{z, x} \text{"non-zero"}$$

z, x "It had better not be -10!"

$$\max \sum_{\text{constant}} \text{"non-zero"} \cdot \log \frac{\beta_2}{\alpha_2} \rightarrow \text{choose } \alpha_2 \text{ as small as possible but } \geq x_i$$

The first part of L, $P(z=z | x=x_i, \theta)$ is a constant w.r.t. θ^* , which we are maximizing. Note that these always have constraints (e.g. $\sum \beta_i = 1$), so we end up with Lagrangians.