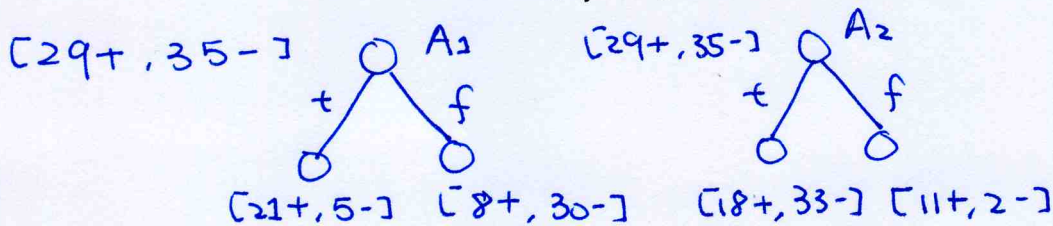


Attribute target Attribute [9+, 5-] root node

Day	Outlook	Temperature	Humidity	Wind	PlayTennis
D1	Sunny	Hot	High	Weak	No
D2	Sunny	Hot	High	Strong	No
D3	Overcast	Hot	High	Weak	Yes
D4	Rain	Mild	High	Weak	Yes
D5	Rain	Cool	Normal	Weak	Yes
D6	Rain	Cool	Normal	Strong	No
D7	Overcast	Cool	Normal	Strong	Yes
D8	Sunny	Mild	High	Weak	No
D9	Sunny	Cool	Normal	Weak	Yes
D10	Rain	Mild	Normal	Weak	Yes
D11	Sunny	Mild	Normal	Strong	Yes
D12	Overcast	Mild	High	Strong	Yes
D13	Overcast	Hot	Normal	Weak	Yes
D14	Rain	Mild	High	Strong	No

leaf node

1. How to decide the root node? Which attribute is better?



how to measure the purity?

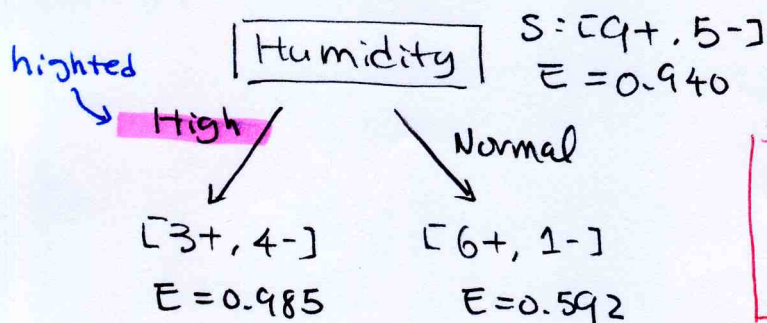
- Entropy: $H(S) = -(P_+ \log P_+ + P_- \log P_-)$ for $\theta, \bar{\theta}$

Generally: $H = -\sum_i P_i \log P_i$

- Information Gain

$$\text{Gain}(S, A) = H(S) - \sum_{v \in \text{Values}(A)} \frac{|S_v|}{|S|} H(S_v)$$

2. Consider the above example:



$$H(S) = -\left(\frac{9}{14} \log \frac{9}{14} + \frac{5}{14} \log \frac{5}{14}\right) = 0.940$$

$$\text{Gain}(S, A = \text{Humidity}) = 0.940 - \left(\frac{3}{14} \times 0.985\right) - \left(\frac{6}{14} \times 0.592\right) = 0.151$$

Similarly: $\text{Gain}(S, \text{Wind}) = 0.048$

$\text{Gain}(S, \text{Humidity}) = 0.151$

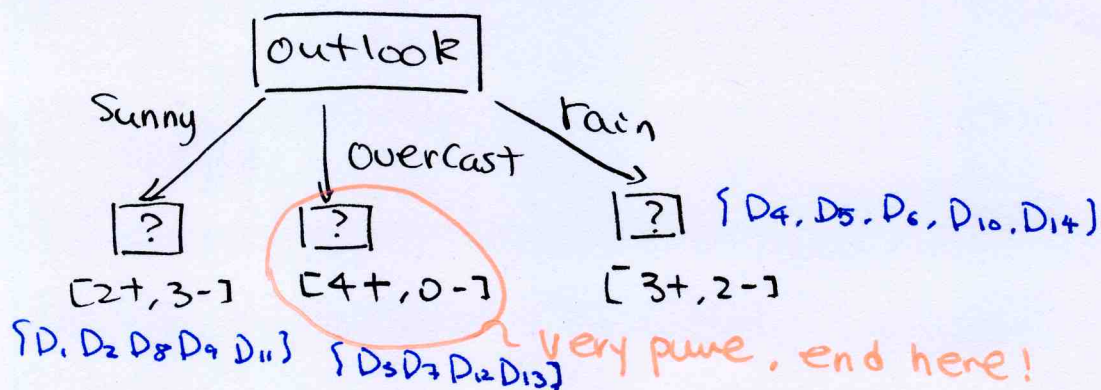
$\text{Gain}(S, \text{Outlook}) = 0.246$

$\text{Gain}(S, \text{temperature}) = 0.029$

maximum Information Gain

So that: the root node for the "playing - tennis" problem is the attribute of "outlook"!

3. Decision Tree Building:

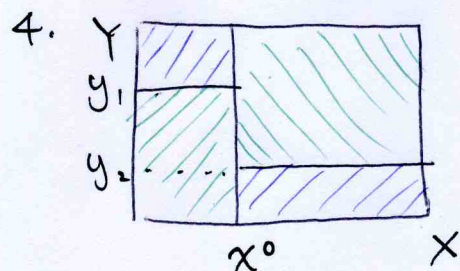


$$H(\text{outlook} = \text{Sunny}) = -\left(\frac{2}{5} \log \frac{2}{5} + \frac{3}{5} \log \frac{3}{5}\right) = 0.97$$

Outlook = Sunny [2+, 3-]

[?] {Humidity, Temperature, Wind}

$$\text{Gain}(\text{Sunny}, \text{Humidity}) = 0.97 - \left(\frac{3}{5}\right) \times 0 - \left(\frac{2}{5}\right) \times 0 = 0.97$$



~ Split of the space.

IF $x < x^0 \wedge y > y_1$ THEN [Blue diagonal lines]
 IF $x < x^0 \wedge y \leq y_1$ THEN [Green diagonal lines]
 IF $x \geq x^0 \wedge y < y_2$ THEN [Blue diagonal lines]
 IF $x \geq x^0 \wedge y \geq y_2$ THEN [Green diagonal lines]

Decision tree is a transparent model comparing to black-box probabilistic models. A tree can be interpreted into a set of decision rules!

Naive Bayes. $P(T|D) = \frac{P(D|T)P(T)}{P(D)} \propto P(D|T)P(T)$
 $= \prod P(D_i|T)P(T)$
 E.g.: $P(T = \text{Yes} | \text{Sunny, hot, normal, weak}) \stackrel{d=1}{=}$

Day	Outlook	Temperature	Humidity	Wind	PlayTennis
D1	Sunny	Hot	High	Weak	No
D2	Sunny	Hot	High	Strong	No
D3	Overcast	Hot	High	Weak	Yes
D4	Rain	Mild	High	Weak	Yes
D5	Rain	Cool	Normal	Weak	Yes
D6	Rain	Cool	Normal	Strong	No
D7	Overcast	Cool	Normal	Strong	Yes
D8	Sunny	Mild	High	Weak	No
D9	Sunny	Cool	Normal	Weak	Yes
D10	Rain	Mild	Normal	Weak	Yes
D11	Sunny	Mild	Normal	Strong	Yes
D12	Overcast	Mild	High	Strong	Yes
D13	Overcast	Hot	Normal	Weak	Yes
D14	Rain	Mild	High	Strong	No

$$P(\text{Sunny} | \text{Yes}) P(\text{hot} | \text{Yes}) P(\text{normal} | \text{Yes})$$

$$P(\text{weak} | \text{Yes}) = \frac{2}{9} \cdot \frac{2}{9} \cdot \frac{6}{9} \cdot \frac{6}{9}$$

$$P(\text{Sunny} | \text{No}) P(\text{hot} | \text{No}) P(\text{normal} | \text{No})$$

$$P(\text{weak} | \text{No}) = \frac{3}{5} \cdot \frac{2}{5} \cdot \frac{1}{5} \cdot \frac{2}{5}$$

$$P(\text{Yes}) = \frac{9}{14}$$

$$P(\text{No}) = \frac{5}{14}$$

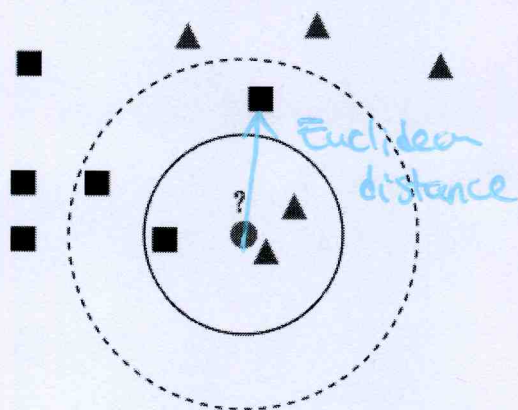
$$P(T = \text{Yes} | X) = \frac{2}{9} \cdot \frac{2}{9} \cdot \frac{6}{9} \cdot \frac{6}{9} \cdot \frac{9}{14} = 0.0141$$

$$P(T = \text{No} | X) = \frac{3}{5} \cdot \frac{2}{5} \cdot \frac{1}{5} \cdot \frac{2}{5} \cdot \frac{5}{14} = 0.0068$$

$P(\text{Yes} | X) > P(\text{No} | X)$

① K-Nearest Neighbour (instance-based learning, lazy learning)

KNN for classification, see the following graph.

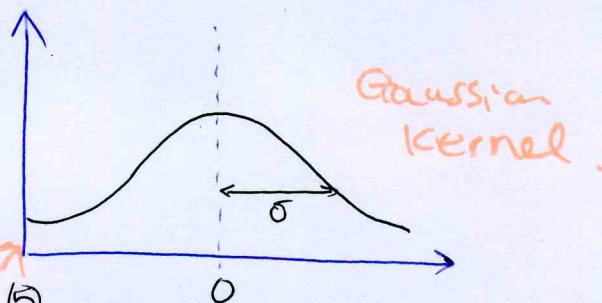


② Radial Basis Function

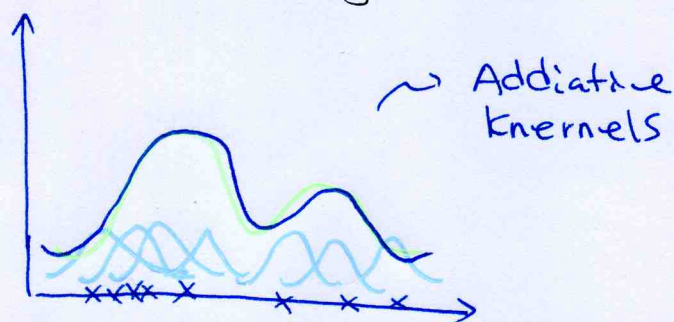
$$\hat{f}(x) = w_0 + \sum_{u=1}^k w_u K_u(d(x_u, x))$$

$K_u(d(x_u, x))$ is a kernel function.

$$K_u(d(x_u, x)) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{1}{2\sigma^2} d(x_u, x)^2\right)$$



③ kernel density estimation.

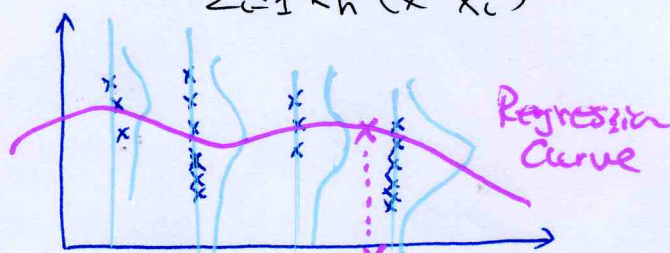


It gives more smooth estimation of the sampling data elements!

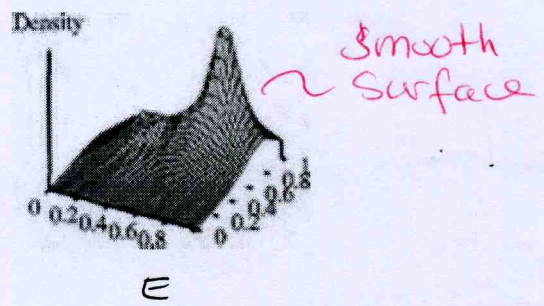
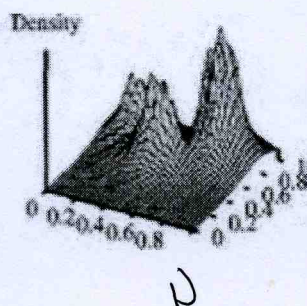
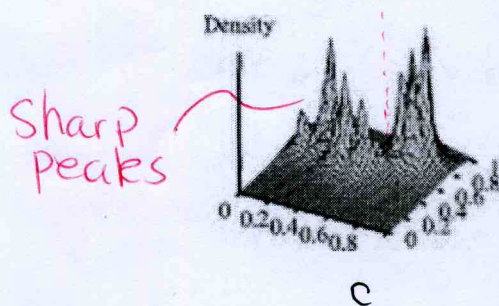
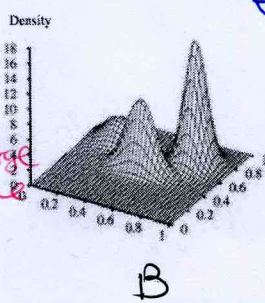
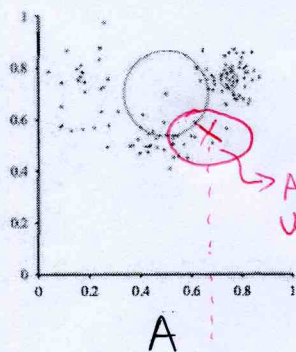
$$\hat{f}_h(x) = \frac{1}{n} \sum_{i=1}^n K_h(x - x_i)$$

⑤ Nadaraya-Watson Kernel Regression: to estimate Conditional Expectation $E(Y|X)$

$$\hat{m}_h(x) = \frac{\sum_{i=1}^n K_h(x - x_i) y_i}{\sum_{i=1}^n K_h(x - x_i)}$$



④ Given the data elements sampled from Fig. B. are shown in Fig. A. Fig. C, D and E are the K-NN density estimation with $k=3, 10$ and 40 .



KNN Regression, ~~and~~ data element x : $\hat{x} = \frac{1}{k} \sum_{i=1}^k (x_i, x)$
where $N(x_i, x)$ is the nearest k data elements given x .

$$\textcircled{1} J(\theta) = \frac{1}{2} \sum_{i=1}^m (h_{\theta}(x^{(i)}) - y^{(i)})^2 \quad \text{Cost Function}$$

$$\theta_j := \theta_j - \alpha \frac{\partial}{\partial \theta_j} J(\theta) \quad \swarrow \text{gradient descent}$$

$$\begin{aligned} \frac{\partial}{\partial \theta_j} J(\theta) &= \frac{\partial}{\partial \theta_j} \frac{1}{2} (h_{\theta}(x^{(i)}) - y^{(i)})^2 \\ &= 2 \cdot \frac{1}{2} (h_{\theta}(x^{(i)}) - y^{(i)}) \frac{\partial}{\partial \theta_j} (h_{\theta}(x^{(i)}) - y^{(i)}) \\ &= (h_{\theta}(x^{(i)}) - y^{(i)}) \frac{\partial}{\partial \theta_j} \left(\sum_{i=0}^n \theta_i x_i^{(i)} - y^{(i)} \right) \\ &= (h_{\theta}(x^{(i)}) - y) x_j^{(i)} \end{aligned}$$

$$\Rightarrow \theta_j = \theta_j - \alpha (h_{\theta}(x^{(i)}) - y) x_j^{(i)}$$

$$= \theta_j + \alpha (y - h_{\theta}(x^{(i)})) x_j^{(i)}$$

$$\begin{cases} h_{\theta}(x) = \theta_0 + \theta_1 x_1 + \theta_2 x_2 \\ h_{\theta}(x) = \sum_i \theta_i x_i = \theta^T x \end{cases}$$

~ LMS Rule
Wiener-Hoff
Learning Rule

② Batch gradient descent

$$\theta_j := \theta_j + \alpha \sum_{i=1}^m (y^{(i)} - h_{\theta}(x^{(i)})) x_j^{(i)}$$

for every $j = 1 \dots n$ (dimension of x)

② Incremental gradient descent

$$\text{for } i = 1 \text{ to } m : \{ \theta_j := \theta_j + \alpha (y - h_{\theta}(x^{(i)})) x_j^{(i)} \}$$

③ Probabilistic Interpretation.

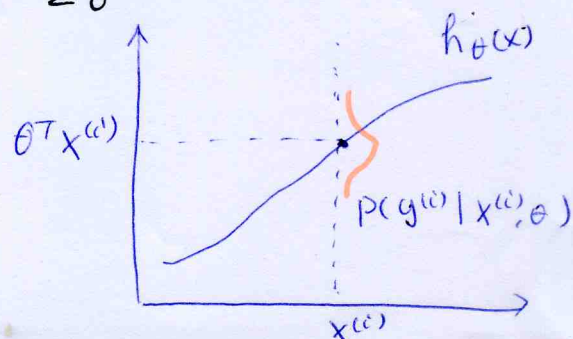
$$y^{(i)} = \theta^T x^{(i)} + \varepsilon^{(i)}$$

$$\varepsilon^{(i)} = (y^{(i)} - \theta^T x^{(i)})$$

if we assume: $P(\varepsilon^{(i)}) \sim N(0, \sigma^2)$

$$P(y^{(i)} | x^{(i)}, \theta) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(y^{(i)} - \theta^T x^{(i)})^2}{2\sigma^2}\right)$$

$$y^{(i)} | x^{(i)}, \theta \sim N(\theta^T x^{(i)}, \sigma)$$



Given the probabilistic interpretation of the LMS.
We can regard the linear regression as maximum likelihood estimation.

$$\mathcal{L}(\theta) = \prod_{i=1}^m P(y^{(i)} | x^{(i)}, \theta)$$

$$y^{(i)} | x^{(i)}, \theta \sim N(\theta^T x^{(i)}, \sigma^2)$$

using log likelihood:

$$\begin{aligned} \ell(\theta) &= \log \mathcal{L}(\theta) = \sum_{i=1}^m \log P(y^{(i)} | x^{(i)}, \theta) \\ &= \sum_{i=1}^m \log \left\{ \frac{1}{\sqrt{2\pi}\sigma} \exp \left\{ -\frac{1}{2\sigma^2} (y^{(i)} - \theta^T x^{(i)})^2 \right\} \right\} \\ &= m \log \frac{1}{\sqrt{2\pi}\sigma} - \frac{1}{2\sigma^2} \sum_{i=1}^m (y^{(i)} - \theta^T x^{(i)})^2 \end{aligned}$$

Maximize log likelihood $\ell(\theta)$ is equivalent to minimize $\sum_{i=1}^m (y^{(i)} - \theta^T x^{(i)})^2$, which is exactly LMS.

⊕ In vector form (or in matrix)

$$\text{Error Term: } E(\theta) = (y - \theta^T x)(y - \theta^T x)^T$$

$$= (y - \theta^T x)(y^T - x^T \theta)$$

$$= yy^T + \theta^T x x^T \theta - \theta^T x y^T - y x^T \theta$$

Some descriptions (dimensions) of variables.

$$y: 1 \times N$$

$$y^T: N \times 1$$

$$\theta: d \times 1$$

$$\theta^T: 1 \times d$$

$$X: d \times N$$

$$X^T: N \times d$$

$$1 \times N, N \times 1$$

$$1 \times d, d \times N, N \times d, d \times 1$$

$$1 \times N, N \times d, d \times 1$$

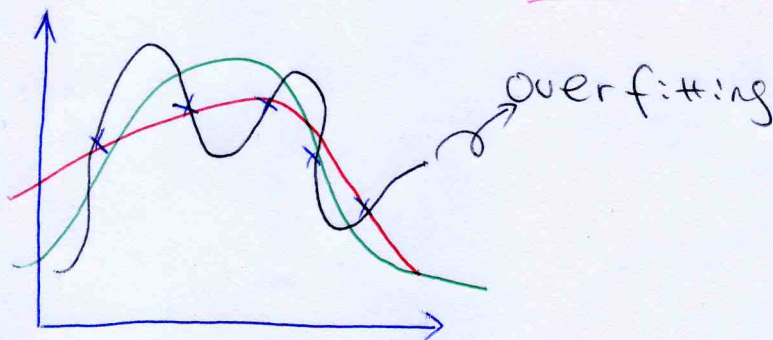
All the terms are scalars.

To take the derivative:

$$\frac{\partial E}{\partial \theta} = 0 \Rightarrow 2Xx^T\theta - 2xy^T = 0$$

$$\Rightarrow \theta = (Xx^T)^{-1} Xy^T$$

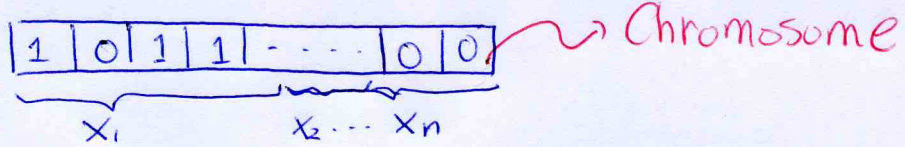
LMS estimation



Genetic Algorithm. ~ Stochastic optimization technique

For a complex function $f(x) = f(x_1, x_2, \dots, x_n)$
and $\{x_1, \dots, x_n\}$ following some constraints $g(x_1, \dots, x_n) = 0$
how to find the maximum value of the given function.

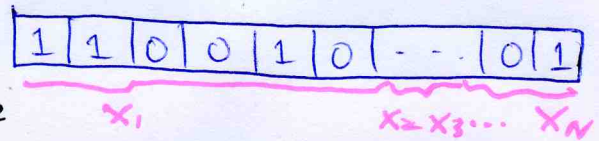
1. Encoding



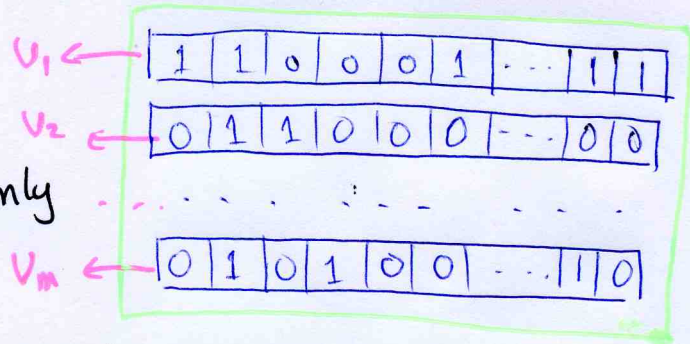
Each variable can be represented by a binary code.

2. Fitness function.

Evaluate a chromosome
by $f(x_1, \dots, x_n)$



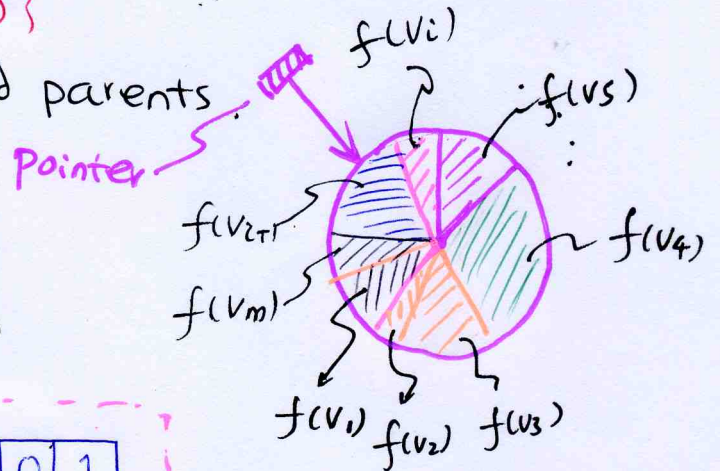
3. A pool of Chromosomes.
which are generated randomly



Each chromosome can be evaluated by the fitness function.
 $\{f(V_1), f(V_2), \dots, f(V_m)\}$

4. Ranking and selecting good parents

The chromosomes with bigger $f(V_i)$ values have better chances to be selected as the parents. E.g:



V_i | 0 | 1 | 1 | 0 | 1 | ... | 0 | 1

V_j | 0 | 0 | 1 | 1 | 1 | ... | 1 | 1

selected parents

V'_i | 0 | 1 | 1 | 0 | 1 | ... | 1 | 1

V'_j | 0 | 0 | 1 | 1 | 1 | ... | 0 | 1

children

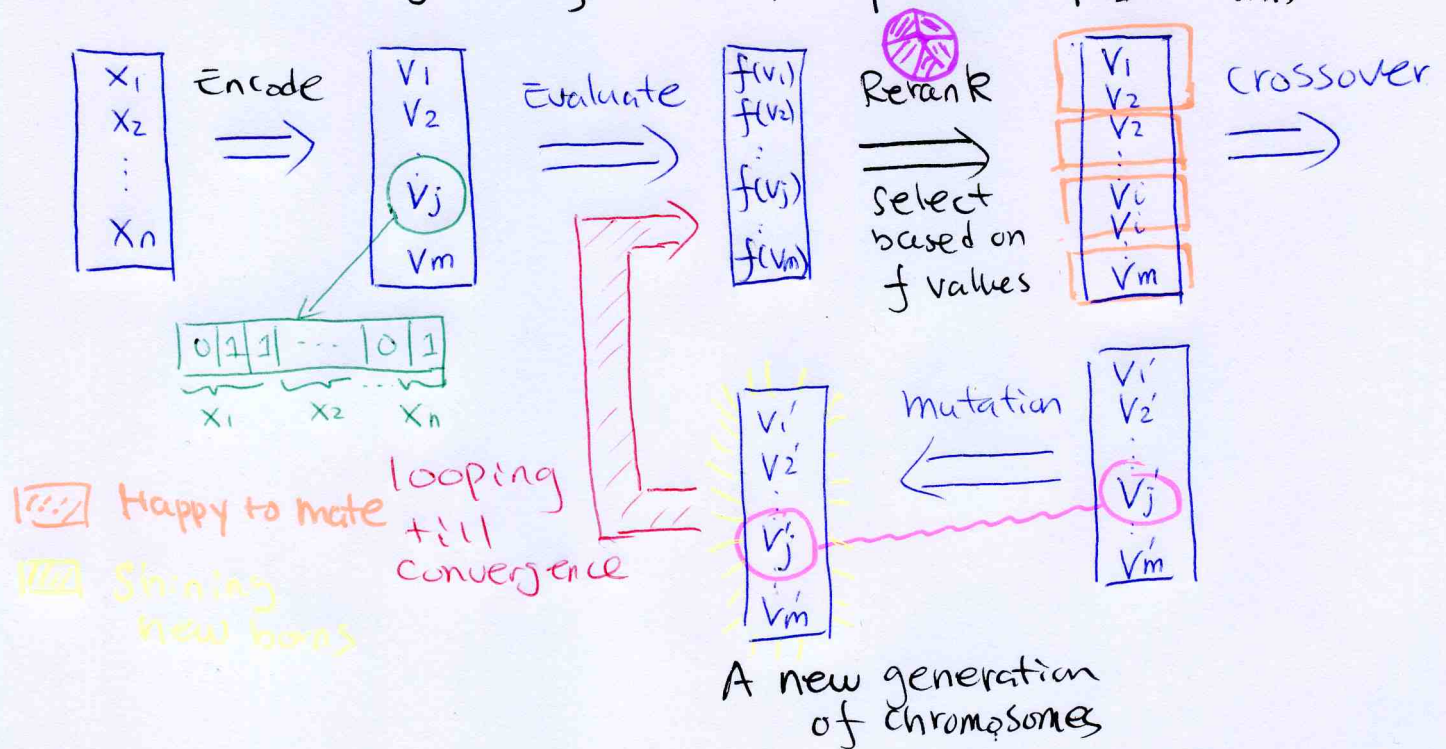
$V_{j'}$ | 0 | 0 | 1 | 0 | 1 | ... | 0 | 1

Crossover

Mutation

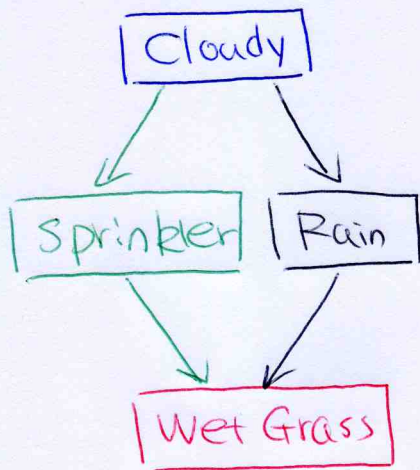
5. GA process.

Initialize a pool of chromosomes $V_1 \dots V_m$, where V_i is encoded by binary codes to represent $\{x_1 \dots x_n\}$



Bayesian Network (variable Elimination method)

prior for Cloudy (C)



$T_1 \quad P(C=F) = P(C=T) = 0.5$

T_2

C	$P(S=F)$	$P(S=T)$
F	0.5	0.5
T	0.9	0.1

T_3

C	$P(R=F)$	$P(R=T)$
F	0.8	0.2
T	0.2	<u>0.8</u>

Reasoning Table for Wet Grass:

$P(R=T|C=T) = 0.8$

T_4

S	R	$P(W=F)$	$P(W=T)$
F	F	1.0	0.0
T	F	0.1	0.9
F	T	0.1	0.9
T	T	0.01	0.99

Suppose all the 3 tables are given from training. How do we do the inference?

E.g. $P(S=T|W=T) = ?$

Using Bayesian Theorem:

① $P(S=T|W=T) = \frac{P(W=T|S=T)P(S=T)}{P(W=T)} = \frac{P(W=T, S=T)}{P(W=T)}$

② $P(W=T, S=T) = \sum_c \sum_r P(C=c, S=T, R=r, W=T)$
 $= \sum_c P(C=c) P(S=T|C=c) \sum_r P(R=r|C=c) \cdot P(W=T|S=T, R=r)$

③ $\sum_r P(R=r|C=c) \cdot P(W=T|S=T, R=r)$
 $= P(R=T|C=c) \cdot P(W=T|S=T, R=T) + P(R=F|C=c) \cdot P(W=T|S=T, R=F)$
 $= P(R=T|C=c) \times 0.99 + P(R=F|C=c) \times 0.9$

② ~ ③ $\Rightarrow$ ④ $P(W=T, S=T) = \sum_c P(C=c) P(S=T|C=c) \times$

④ $P(W=T, S=T) = P(C=T) P(S=T|C=T) (0.99 P(R=T|C=T) + 0.9 P(R=F|C=T))$
 $+ \cancel{0.5} P(C=F) P(S=T|C=F) (0.99 P(R=T|C=F) + 0.9 P(R=F|C=F))$
 $= 0.5 \times 0.1 \times (0.99 \times 0.8 + 0.9 \times 0.2) + 0.5 \times 0.5 \times (0.99 \times 0.2 + 0.9 \times 0.8)$
 $= 0.05 \times 0.972 + 0.25 \times 0.918 = 0.2781$

⑤ $P(W=T) = \sum_c \sum_r \sum_s P(C=c, S=s, R=r, W=T)$
 $= \sum_c P(C=c) \underbrace{\sum_s P(S=s|C=c)}_{\{C\}} \underbrace{\sum_r P(R=r|C=c) P(W=T|S=s, R=r)}_{\{C, S\}}$

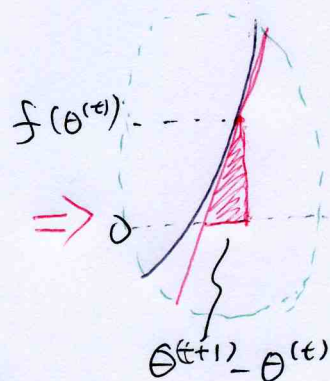
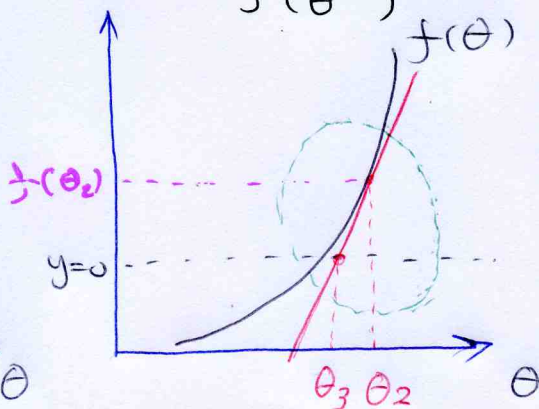
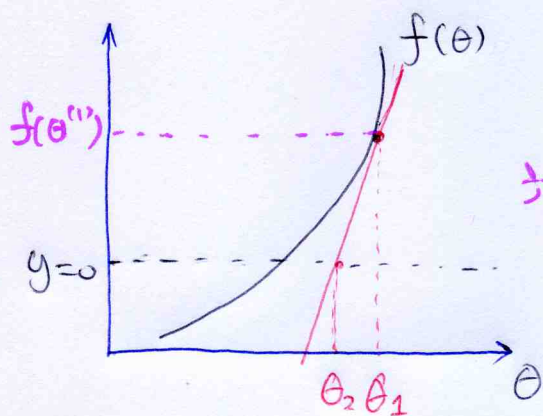
Newton's Method for maximizing likelihood.

① Newton's method for finding a zero of a function.

$f: \mathbb{R} \rightarrow \mathbb{R}$, and wish to find the value of θ , so that $f(\theta) = 0$, $\theta \in \mathbb{R}$ is a real number. Then θ is updated based on:

$$\theta \leftarrow \theta - \frac{f(\theta)}{f'(\theta)}$$

or formally:
$$\theta^{(t+1)} \leftarrow \theta^{(t)} - \frac{f(\theta^{(t)})}{f'(\theta^{(t)})}$$



② In maximizing likelihood.

$\frac{\partial \ell}{\partial \theta} = 0$, so that
$$\theta^{(t+1)} \leftarrow \theta^{(t)} - \frac{\ell'(\theta^{(t)})}{\ell''(\theta^{(t)})}$$

$$\frac{\theta^{(t+1)} - \theta^{(t)}}{f(\theta^{(t)})} = \frac{1}{f'(\theta^{(t)})}$$

To generalize the Newton's method.

(Newton - Raphson Method) is given by:

$$\vec{\theta}^{(t+1)} \leftarrow \vec{\theta}^{(t)} - H^{-1} \nabla_{\theta} \ell(\theta)$$

where H is the Hessian matrix

$$H_{ij} = \frac{\partial^2 \ell(\theta)}{\partial \theta_i \partial \theta_j}$$

$$\nabla_{\theta} \ell(\theta) = \frac{\partial \ell(\theta)}{\partial \theta_i} \text{ for } i=1 \dots d$$

③ Newton method is good for convex function maximization.

$f: \mathbb{R}^n \Rightarrow \mathbb{R}$ is convex if

$$f(\theta x + (1-\theta)y) \leq \theta f(x) + (1-\theta)f(y) \quad (0 \leq \theta \leq 1)$$

