

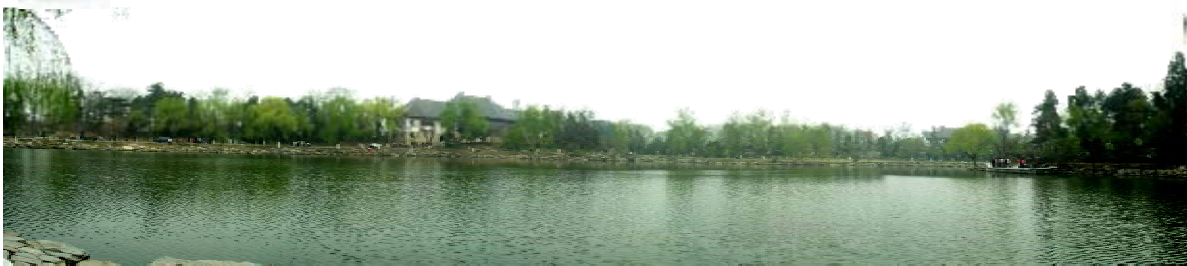


Introduction to Computational Biology

《计算生物学导论》

第七章

计算生物学的 随机过程方法



Peking University

Introduction to Computational Biology



A mainstream topic in computational biology is the problem of sequence annotation: given a sequence of DNA/RNA or protein, ..., etc., we want to identify “interesting” elements

Examples:

- DNA/RNA: genes, promoters, splicing signals, binding sites, etc
- Protein: coiled-coil domains, transmembrane domains, signal peptides, phosphorylation sites, etc
- Generally: homologs, etc





The sequence of many of these interesting elements can be characterized statistically, so we are interested in modeling them.

By modeling, we mean find statistical models then can:

- Accurately describe the observed elements of provided sequences;
- Accurately predict the presence of particular elements in new or un-annotated sequences;
- If possible, be readily interpretable and provide some insight into the actual biological process involved (*i.e.* not a black box).



随机过程 (Random process, or Stochastic process)

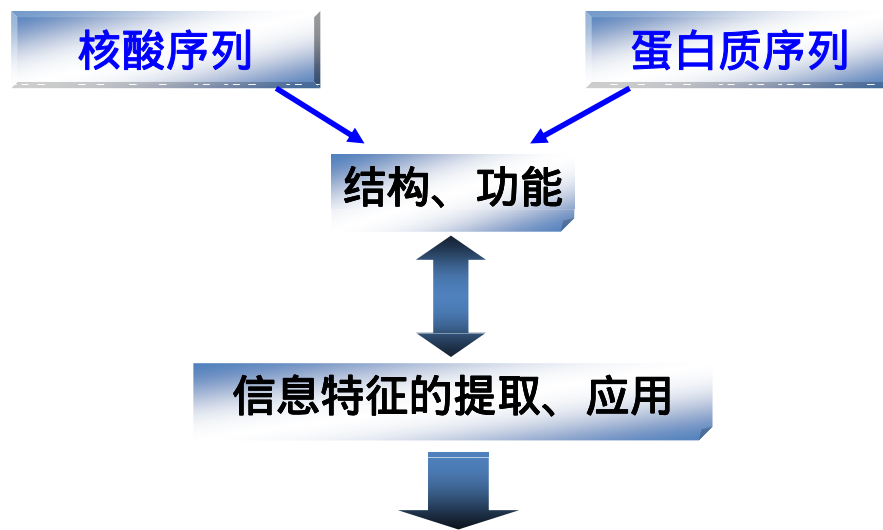
A collection of random variables, representing the evolution of some a system of random values over time, which is usually just a kind of mathematical expression.

- A **stochastic process** X is a collection of random variables $\{X_t, t \in T\}$
- The variable t is often interpreted as time, and $X(t)$ is the state of the process at time t .
- A realization of X is called a **sample path**.
- There are a number of special cases that are of high interest in computational biology, in particular **Markov processes**.





§ 7.1 生物序列的Markov过程模型



方法：建立符号（核苷酸、氨基酸）序列的概率模型

AGGGAGGGTAGGACTAGATGTTTTGGA



生物序列分析的两类问题

(1) 判别问题

Does the sequence belong to a particular structure class (family)?

(2) 信息结构的识别

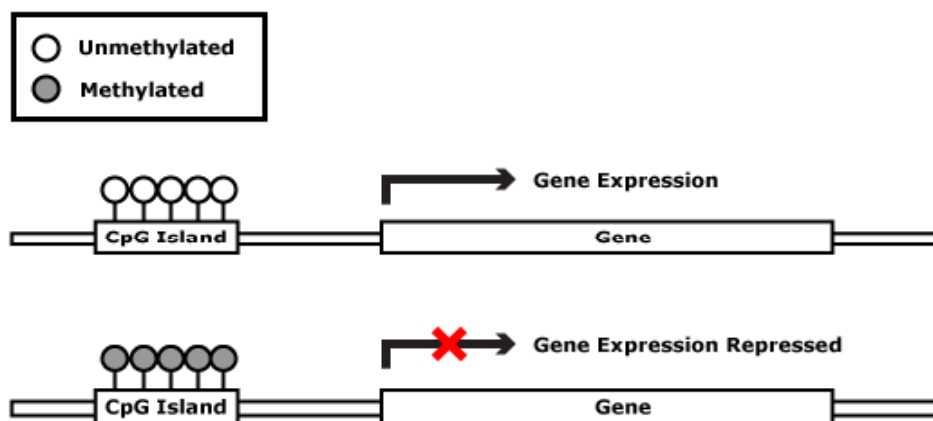
Assuming the sequence includes some structure classes (families), what can we say about its internal structure?



例子：人类基因组中与启动子相关的CpG岛信号

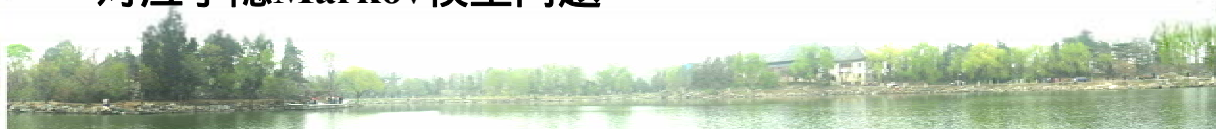
在人类基因组的许多基因的转录起始区域中，二核苷酸CG出现的频率往往要高于基因组其它的区域，因此，CG的高含量区（通常长达 $\sim 10^3$ 碱基）可能意味着转录启动子的存在，这种“CG”高含量区被称为CpG岛（CpG islands）。对CpG岛的识别，有助于转录起始信号的识别。

【**CpG岛（CpG islands）**：指DNA上一个区域，此区域含有大量相连的胞嘧啶（C）、鸟嘌呤（G）以及使两者相连的磷酸酯键（p）。哺乳类基因中的启动子上，含有约40%的CpG岛（人类约70%）。一般CpG岛的长度约300到3000bp。常用的正式定义是指一个至少含有200bp的区域，其中GC所占比例超过50%，且CpG的观察值 / 预测值比例必须高于0.6。此部分的CpG岛与基因相连，可用来作为限制酶的识别位置。】



两个问题：

- (1) 给定一段DNA序列片段，判别它是否CpG岛？
对应于Markov模型问题
- (2) 给定一段DNA序列，识别其中的CpG岛？
对应于隐Markov模型问题





Andrei A Markov was a graduate of Saint Petersburg University (1878), where he began a professor in 1886. His early work was mainly in number theory and analysis, continued fractions, limits of integrals, approximation theory and the convergence of series.

After 1900 Markov applied the method of continued fractions, pioneered by his teacher Pafnuty Chebyshev, to probability theory. He proved the central limit theorem under fairly general assumptions.

Markov is particularly remembered for his study of Markov chains, sequences of random variables in which the future variable is determined by the present variable but is independent of the way in which the present state arose from its predecessors. This work launched the theory of stochastic processes.

Andrei Andreyevich Markov
(1856~1922)



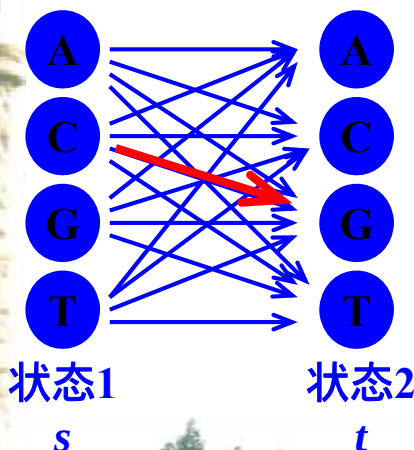
A. A. Markov. (1886)



1) 符号序列的Markov模型

...ACGCCGATCCTCGCGACGACGGTGCGACGTCGTCCG...

例：CpG岛区域
的双核苷酸



序列模型的假设：

一种（符号）状态转移到另一种（符号）状态时，后一状态仅取决于前一状态。



1阶Markov过程（Markov chain）

这种假设丢掉了许多信息，但是得到了复杂序列的简化的概率模型。

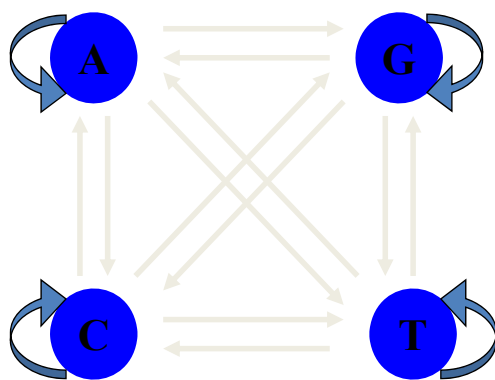


定 义

Markov过程：从一种状态转移到另一种状态时，过程仅取决于前面 n 种状态，是一种有序 n 模型。 n 是影响下一个状态选择的的状态数，由此定义的称之为 **n 阶Markov过程**。

最简单的Markov过程是一阶过程，状态的选择完全取决于前一状态，这种选择是依照概率来选择的。

注意：状态的选择是概率的，而非确定的。故Markov过程本质上是一种随机过程。



例：某段DNA序列的1阶Markov过程模型 ($s \rightarrow t$)

	t A	t C	t G	t T
s A	0.18	0.27	0.42	0.12
s C	0.17	0.37	0.27	0.19
s G	0.16	0.34	0.38	0.13
s T	0.08	0.36	0.38	0.18

定义：**Markov过程**

由下列参数构成：

(1) **状态 (state)**：A、C、G、T

(2) **状态转移概率 (state transition probability) 矩阵**：AA, AC, ... 的概率大小，记为 a_{st}



2) DNA序列的1阶Markov模型

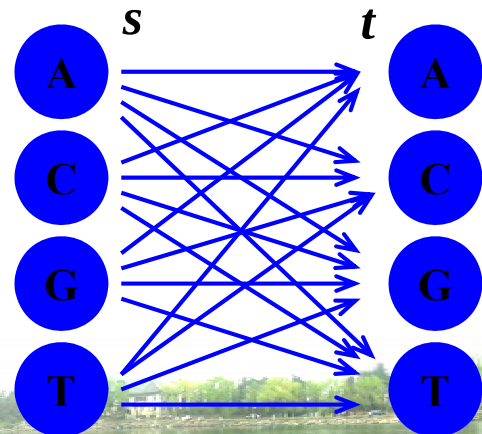
假定DNA序列长度为 L ，序列 x 写成 $x_1x_2\dots x_L$ 。

状态转移概率为 a_{st}

$$a_{st} = P(x_i = t \mid x_{i-1} = s) \quad i = 2, 3, \dots, L$$

a_{st} 表示碱基 s 后面紧邻碱基 t 的概率，或者状态 s 后紧邻状态 t 的概率。

(a_{st} 标志状态 s 转移到状态 t 的箭头联系的强度。)



运用乘法公式，可以写出序列 $x_1x_2\dots x_L$ 的一般意义下的概率模型表达式：

$$\begin{aligned} P(x) &= P(x_L, x_{L-1}, \dots, x_1) \\ &= P(x_L \mid x_{L-1}, \dots, x_1) P(x_{L-1} \mid x_{L-2}, \dots, x_1) \dots P(x_1) \end{aligned}$$

1阶Markov过程的特点是碱基 x_i 只取决于它前面 x_{i-1} ，而非整个序列 x 。

亦即：

$$P(x_i \mid x_{i-1}, \dots, x_1) = P(x_i \mid x_{i-1}) = a_{x_{i-1}x_i}$$



因此，

$$\begin{aligned} P(x) &= P(x_L, x_{L-1}, \dots, x_1) \\ &= P(x_L | x_{L-1}, \dots, x_1) P(x_{L-1} | x_{L-2}, \dots, x_1) \dots P(x_1) \end{aligned}$$

可写成：

$$\begin{aligned} P(x) &= P(x_L | x_{L-1}) P(x_{L-1} | x_{L-2}) \dots P(x_2 | x_1) P(x_1) \\ &= P(x_1) \prod_{i=2}^L a_{x_{i-1}x_i} \end{aligned}$$

$$P(x) = P(x_1) \prod_{i=2}^L a_{x_{i-1}x_i}$$



练习

证明：所有长度为L的序列的概率总和为1，即

$$\begin{aligned} \sum_{\{x\}} P(x) &= \sum_{x_1} \sum_{x_2} \dots \sum_{x_L} P(x_1) \prod_{i=2}^L a_{x_{i-1}x_i} \\ &= 1 \end{aligned}$$



3) DNA序列的起始状态和结束状态

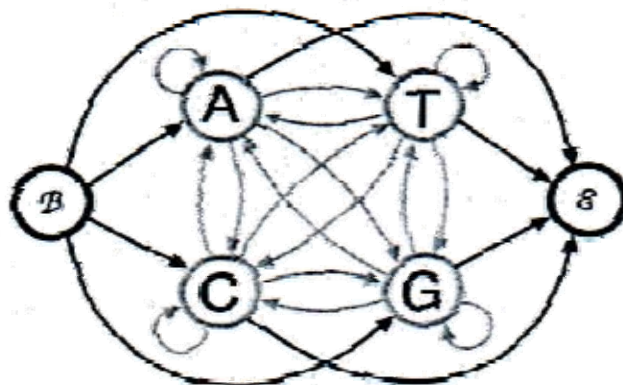
为描述序列起始和结束的状态转移，在A、C、G、T 4个状态之外，人为添加定义2个状态符号B（起始）、E（结束）。

对于序列的起始，定义 $x_0=B$ ，对应 $x_0 \rightarrow x_1$ 的转移概率为：

$$a_{Bs} = P(x_1 = s)$$

对于序列的结束，对应 $x_L \rightarrow E$ 的转移概率为：

$$a_{tE} = P(E | x_L = t)$$



4) CpG islands判别的应用

$$P(x) = P(x_L | x_{L-1})P(x_{L-1} | x_{L-2}) \dots P(x_2 | x_1)P(x_1)$$

$$= P(x_1) \prod_{i=2}^L a_{x_{i-1}x_i}$$

序列x属于某一
Markov过程的似然率

衡量序列x是否满足
某一Markov模型的度量

判别分析问题



已知48条人类DNA序列（全长60kb），每条都含有确认的CpG islands位置、长度注释信息。据此构造训练集：

正方训练集
(Positive)

CpG islands区的DNA序列

负方训练集
(Negative)

CpG islands区之外的DNA序列

分别从positive训练集和negative训练集中统计出双核苷酸 st 的频率，再计算概率转移矩阵：

$$a_{st}^{+} = \frac{c_{st}^{+}}{\sum_{t'} c_{st'}^{+}}$$

$$a_{st}^{-} = \frac{c_{st}^{-}}{\sum_{t'} c_{st'}^{-}}$$



a_{st}^{+}

+	t A	t C	t G	t T
s A	0.18	0.27	0.43	0.12
s C	0.17	0.37	0.27	0.19
s G	0.16	0.34	0.38	0.13
s T	0.08	0.36	0.38	0.18

a_{st}^{-}

-	t A	t C	t G	t T
s A	0.30	0.21	0.29	0.21
s C	0.32	0.30	0.08	0.30
s G	0.25	0.25	0.30	0.21
s T	0.18	0.24	0.29	0.29



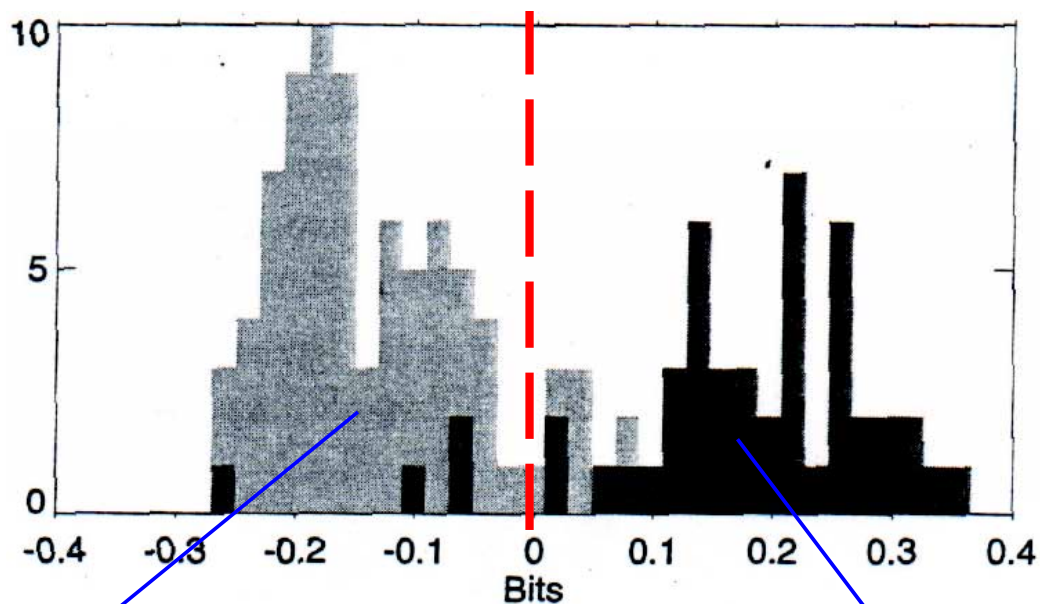
对于任一段DNA序列片段，计算下列分值：

$$S(x) = \log \frac{P(x | \text{model } +)}{P(x | \text{model } -)} = \sum_{i=1}^L \log \frac{a_{x_{i-1}x_i}^+}{a_{x_{i-1}x_i}^-} = \sum_{i=1}^L \beta_{x_{i-1}x_i}$$

β	t A	t C	t G	t T
s A	-0.74	0.42	0.58	-0.80
s C	-0.91	0.30	1.81	-0.69
s G	-0.62	0.46	0.33	-0.73
s T	-1.17	0.57	0.39	-0.68

作长度L的归一处理：

$$S(x) = \frac{1}{L} \sum_{i=1}^L \beta_{x_{i-1}x_i}$$



非CpG岛

CpG岛

据此，可以对某一DNA片段构造判别规则，判定是否CpG岛。



思考：

根据前面的训练集及训练出的参数，如何设计在某一长的DNA序列上识别（注释）CpG岛的大致位置的方法？

...AACCTTCGCGGCCAGGCCGGCTGGTCGTGCGCCATATGTAAGGCC

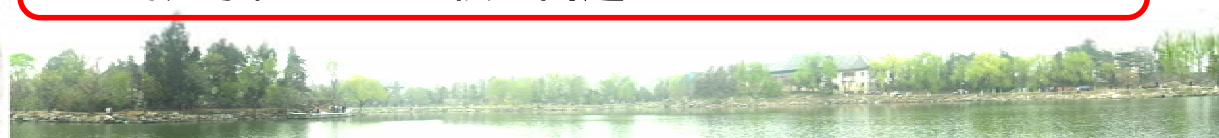


例子：人类基因组中与启动子相关的CpG islands信号

在人类基因组的许多基因的转录起始区域，二核苷酸“CG”出现的频率往往要高于基因组其它的区域，因此，“CG”的高含量区（通常长达数百～数千碱基）可能意味着转录启动子的存在，这种“CG”高含量区被称为CpG岛（CpG islands）。对CpG岛的识别，有助于转录起始信号的识别。

两个问题：

- (1) 给定一段DNA序列片段，判别它是否CpG islands？
对应于Markov模型问题
- (2) 给定一段DNA序列，识别其中的CpG islands？
对应于隐Markov模型问题





§ 7.2 生物序列的隐Markov模型

(隐) Markov模型
(Hidden) Markov model

语音识别
(Speech recognition)

光字符识别
(Optical character recognition)

生物序列分析
(Biological sequence analysis)

生物特征识别
(Biometrics)

- (1) 序列比较与搜寻 (尤其是多序列比对)
- (2) 基因识别、预测 (包括DNA编码与非编码区的识别、真核基因剪接位点信号识别、非编码区的转录调控信号识别、跨膜蛋白的识别.....)
- (3) 蛋白质二级结构、家族、超家族预测、分类
-



语音识别的HMM方法

根据记录的语音波形，分解成若干段10~20毫秒长的声学片段(帧)，每一帧对应于一个符号(通过HMM训练算法获得的声学特征片段)继而转换成符号序列。再根据已经建立的声学模型对该符号序列进行识别，判断对应的音素、音节、单词。





语音识别的基本问题

(1) 语音特征提取：从语音波形中提取出随时间变化的语音特征序列。（模式的构造）

(2) 声学模型与模式匹配（识别算法）：对获取的语音特征通过学习算法产生声学模型。在识别时将输入的语音特征同声学模型（模式）进行匹配与比较，得到最佳的识别结果。

(3) 语言模型与语言处理：语言模型包括由识别语音命令构成的语法网络或由统计方法构成的语言模型，语言处理可以进行语法、语义分析。对小词表语音识别系统，往往不需要语言处理部分。



语音识别的困难

(1) 真实语音发声的多样性 → 庞大的语音数据库和语料数据库

(2) 真实语音状态驻留时间的多样性 → 复杂的随机过程



生物序列分析的挑战

(1) 分析、识别分子序列（4种核苷酸符号或20种氨基酸符号）的生物学意义；共同之处：“Based on a sequence of symbols from some alphabet, find out what the sequence represents.”

(2) 分子序列的复杂性：

结构的多样性

特征序列的保守与突变并存



例子：人类基因组中与启动子相关的CpG岛信号

(1) 给定一段DNA序列片段，判别它是否CpG岛？

对应于Markov过程问题

(2) 给定一段长DNA序列，识别其中的CpG岛？

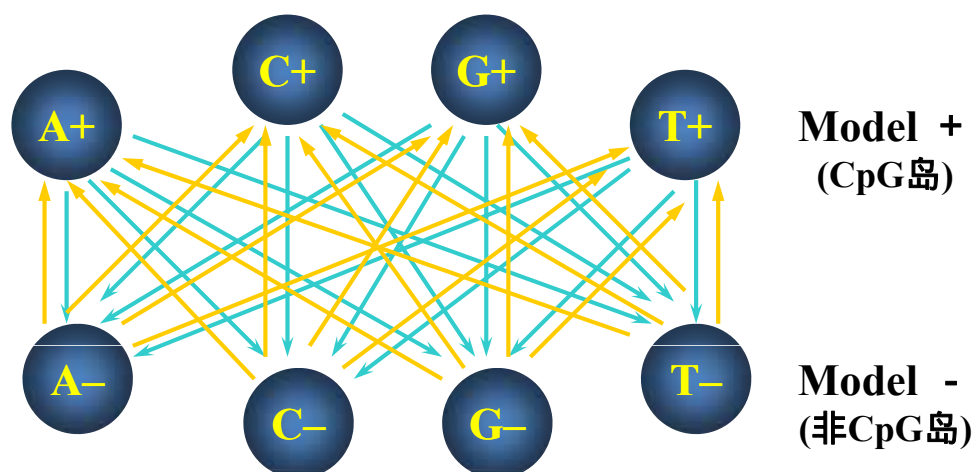
对应于隐Markov模型问题



构造模型，用于刻画较长DNA序列中的非均匀结构
(CpG岛+非CpG岛区)



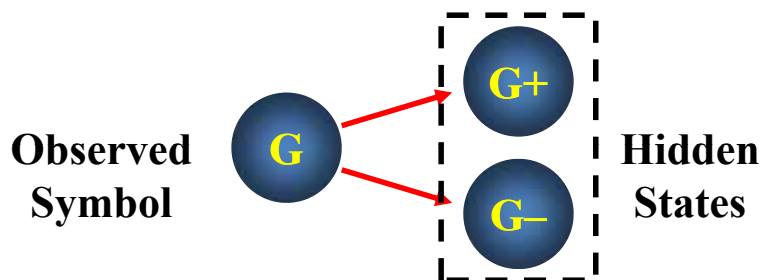
- 构造描述CpG islands的“Model +”和描述非CpG islands的“Model -”混合在一条DNA序列中的Markov过程
- 增加CpG islands与非CpG islands交界的状态转移



刻画含CpG岛结构的DNA序列的HMM模型

(注：除图中所示的状态转移外，“Model +”和“Model -”内的状态转移的定义仍然保留，但图中未标示)





观察的符号与该符号状态的关系

Markov过程： 一对一的关系
HMM: 一对多的关系



1. HMM的定义

观察
(符号) 序列 $x \dots$ $x_{i-1} \ x_i$ **C G G C G A T A** $\dots$ Observed

状态序列 $\pi \dots$ $\pi_{i-1} \ \pi_i$ **C- G+ G+ C+ G+ A+ T- A-** $\dots$ Hidden

Notes：观察（符号）序列 x 和状态序列 π ，都用Markov过程来描述。状态序列 π 也称为路径（path）



状态序列 π 也是Markov过程，故定义状态转移概率 a_{kl} 来刻画状态序列 π ：

$$a_{kl} = P(\pi_i = l \mid \pi_{i-1} = k)$$

同样，人为添加状态序列的起始和结束状态：

$$a_{0l} = P(\pi_1 = l \mid \pi_0 = START) = P(\pi_1 = l)$$

$$a_{k0} = P(\pi_0 = END \mid \pi_L = k) = P(\pi_L = k)$$

建立了两个Markov链：符号序列 $\{x_i\}$ 、状态序列 $\{\pi_i\}$

二者独立



为描述符号序列 $\{x_i\}$ 和状态序列 $\{\pi_i\}$ 的关系，定义参数输出概率 $e_k(b)$ (emission probability)为：

$$e_k(b) = P(x_i = b \mid \pi_i = k)$$

表示在状态 π_i 为 k 时，该状态导致产生符号 b （可以是单个符号，也可以是字符串）的概率。

Notes：状态转移概率 a_{kl} 和输出概率 $e_k(b)$ 的定义都是一种条件概率。



例1：GpC island的HMM模型参数

输出（观察）符号：{A, C, G, T}

状态：{+, -} or {A+, A-, C+, C-, G+, G-, T+, T-}

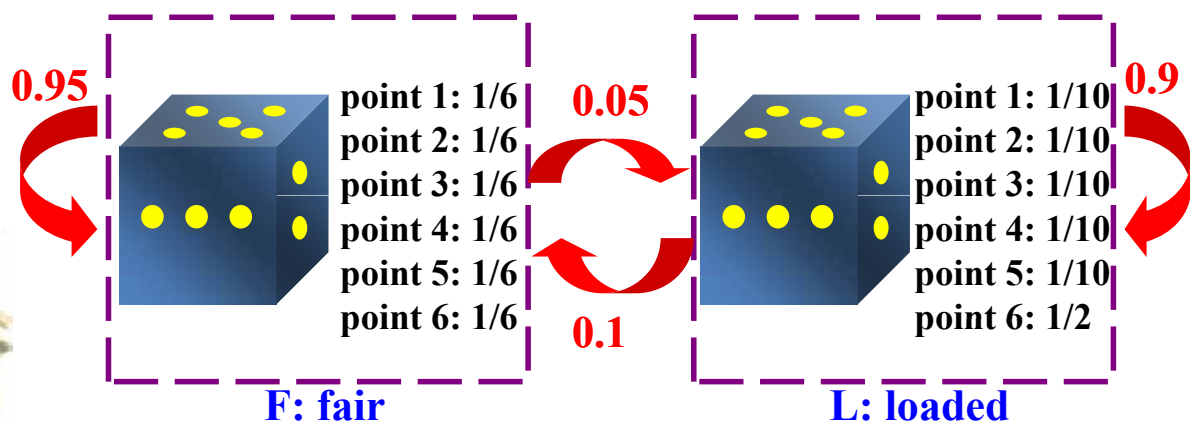
状态转移概率（ p 为GpC岛占全部DNA序列的比例，且 $p+q=1$ ）：

$\pi_i \backslash \pi_{i+1}$	A ⁺	C ⁺	G ⁺	T ⁺	A ⁻	C ⁻	G ⁻	T ⁻
A ⁺	0.180 p	0.274 p	0.426 p	0.120 p	$\frac{1-p}{4}$	$\frac{1-p}{4}$	$\frac{1-p}{4}$	$\frac{1-p}{4}$
C ⁺	0.171 p	0.368 p	0.274 p	0.188 p	$\frac{1-p}{4}$	$\frac{1-p}{4}$	$\frac{1-p}{4}$	$\frac{1-p}{4}$
G ⁺	0.161 p	0.339 p	0.375 p	0.125 p	$\frac{1-p}{4}$	$\frac{1-p}{4}$	$\frac{1-p}{4}$	$\frac{1-p}{4}$
T ⁺	0.079 p	0.355 p	0.384 p	0.182 p	$\frac{1-p}{4}$	$\frac{1-p}{4}$	$\frac{1-p}{4}$	$\frac{1-p}{4}$
A ⁻	$\frac{1-q}{4}$	$\frac{1-q}{4}$	$\frac{1-q}{4}$	$\frac{1-q}{4}$	0.300 q	0.205 q	0.285 q	0.210 q
C ⁻	$\frac{1-q}{4}$	$\frac{1-q}{4}$	$\frac{1-q}{4}$	$\frac{1-q}{4}$	0.322 q	0.298 q	0.078 q	0.302 q
G ⁻	$\frac{1-q}{4}$	$\frac{1-q}{4}$	$\frac{1-q}{4}$	$\frac{1-q}{4}$	0.248 q	0.246 q	0.298 q	0.208 q
T ⁻	$\frac{1-q}{4}$	$\frac{1-q}{4}$	$\frac{1-q}{4}$	$\frac{1-q}{4}$	0.177 q	0.239 q	0.292 q	0.292 q

输出概率： $e_k(b) = P(x_i = b | \pi_i = k)$ 具体值由上述矩阵给出



例2：某赌场骰子的HMM模型参数



输出（观察）符号：{1, 2, 3, 4, 5, 6}

状态：{F, L}

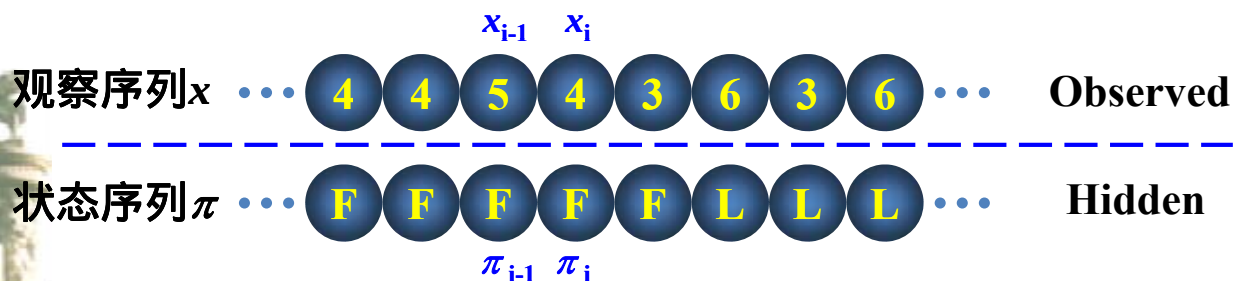
$$a_{FF} = 0.95 \quad a_{FL} = 0.05 \quad a_{LL} = 0.90 \quad a_{LF} = 0.10$$

$$e_F(1) = e_F(2) = e_F(3) = e_F(4) = e_F(5) = e_F(6) = 1/6$$

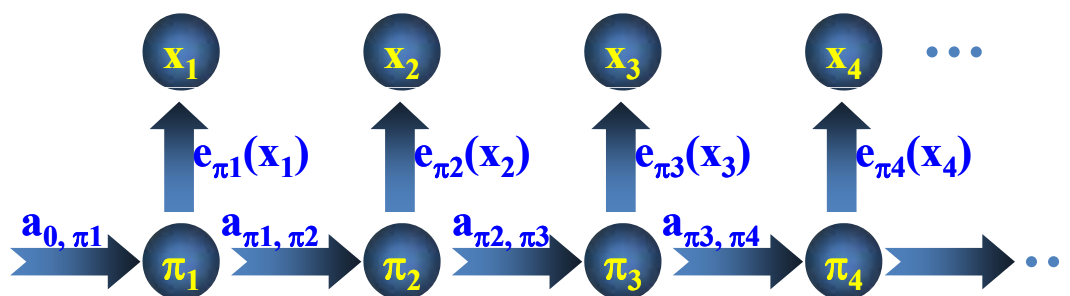
$$e_L(1) = e_L(2) = e_L(3) = e_L(4) = e_L(5) = 0.1$$

$$e_L(6) = 0.5$$





HMM产生观察符号序列的机制：

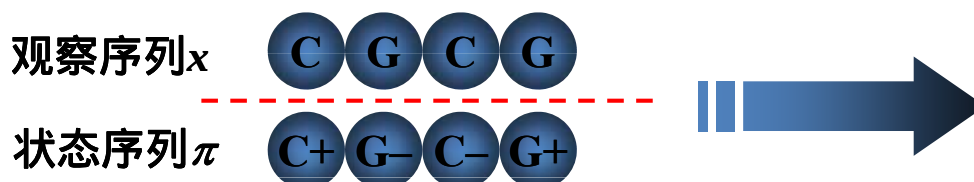




对于观察序列 x 和状态路径 π ，可计算其联合概率：

$$P(x, \pi) = a_{0\pi_1} \prod_{i=1}^L e_{\pi_i}(x_i) a_{\pi_i\pi_{i+1}}$$

例：



由该状态序列产生该观察序列的概率为：

$$a_{0,C_+} \times (1 \times a_{C_+G_-}) \times (1 \times a_{G_-C_-}) \times (1 \times a_{C_-G_+}) \times (1 \times a_{G_+,0})$$



One might hope that knowledge of a good statistical model could suggest a physical model, *i.e.* that the hidden Markov model states actually correspond to physical states of the biological system, but this not required.

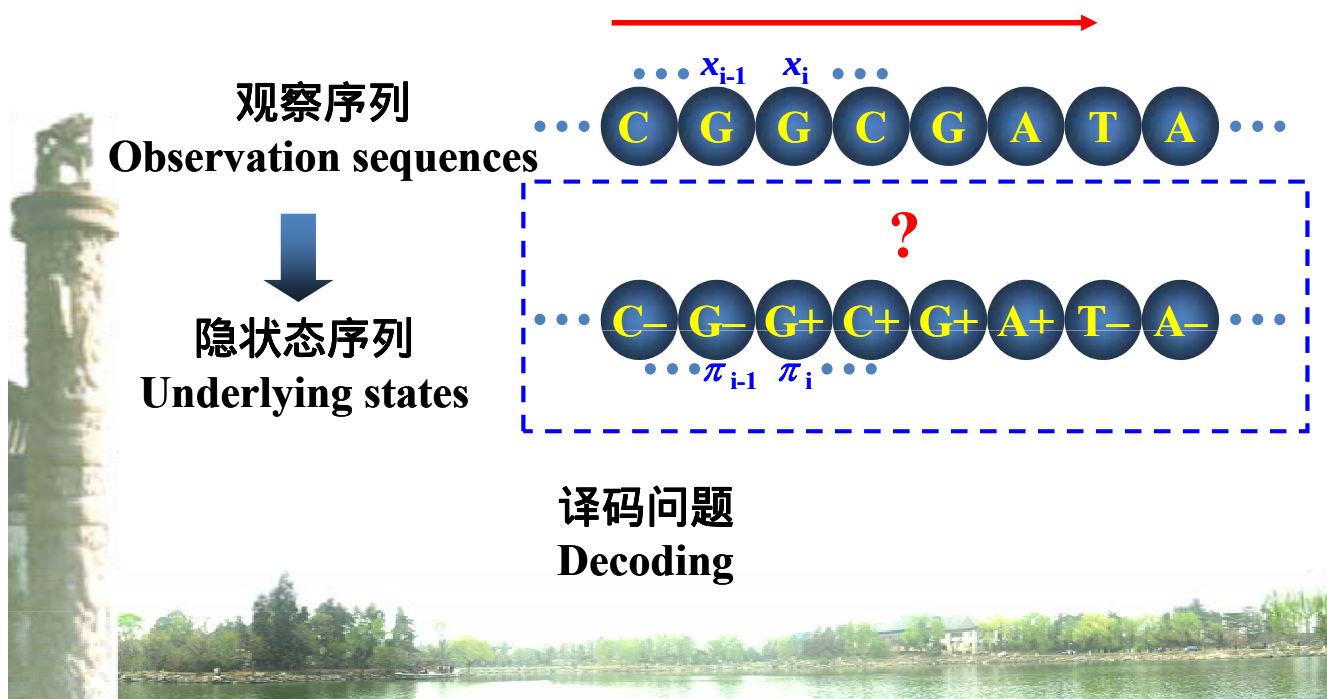
——《Computational Molecular Biology: An Introduction》

本质上说，所有模型都是错误的，但有些是有用的。

—— George Box



2. Viterbi算法——最大可能路径



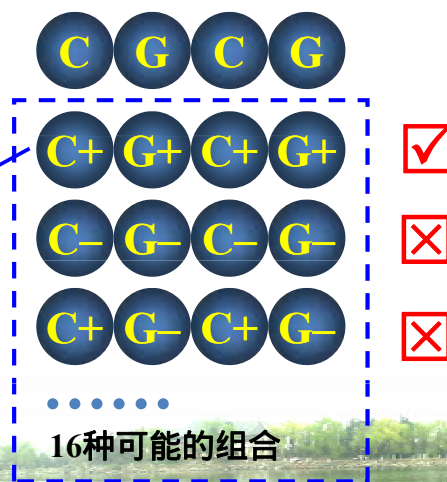
预测隐状态序列的准则

选择最大可能状态路径（Most probable state path），即观察序列 $x_1 \dots x_L$ 发生概率最高的路径，作为寻求的隐状态路径。

$$\pi^* = \arg \max_{\pi} P(x, \pi) = \arg \max_{\pi} P(\pi | x)$$

最大可能路径

$$P(x, \pi) = a_{0\pi_1} \prod_{i=1}^L e_{\pi_i}(x_i) a_{\pi_i \pi_{i+1}}$$





递归方法与动态规划

根据递归 (recursive) 方法来求得最大可能路径。

定义 $v_k(i)$ 为最大可能路径延伸到第 i 步 (即产生前 i 个字串 $x_1 x_2 \dots x_i$)，且序列第 i 步为状态 k 的概率。对各状态 k 的 $v_k(i)$ 为：

$$v_k(i) = \max_{\{\pi | \pi_i = k\}} P(x_1, x_2, \dots, x_i, \pi)$$

则沿状态序列，由动态规划原理得到第 $i+1$ 步的状态为 l 的概率：

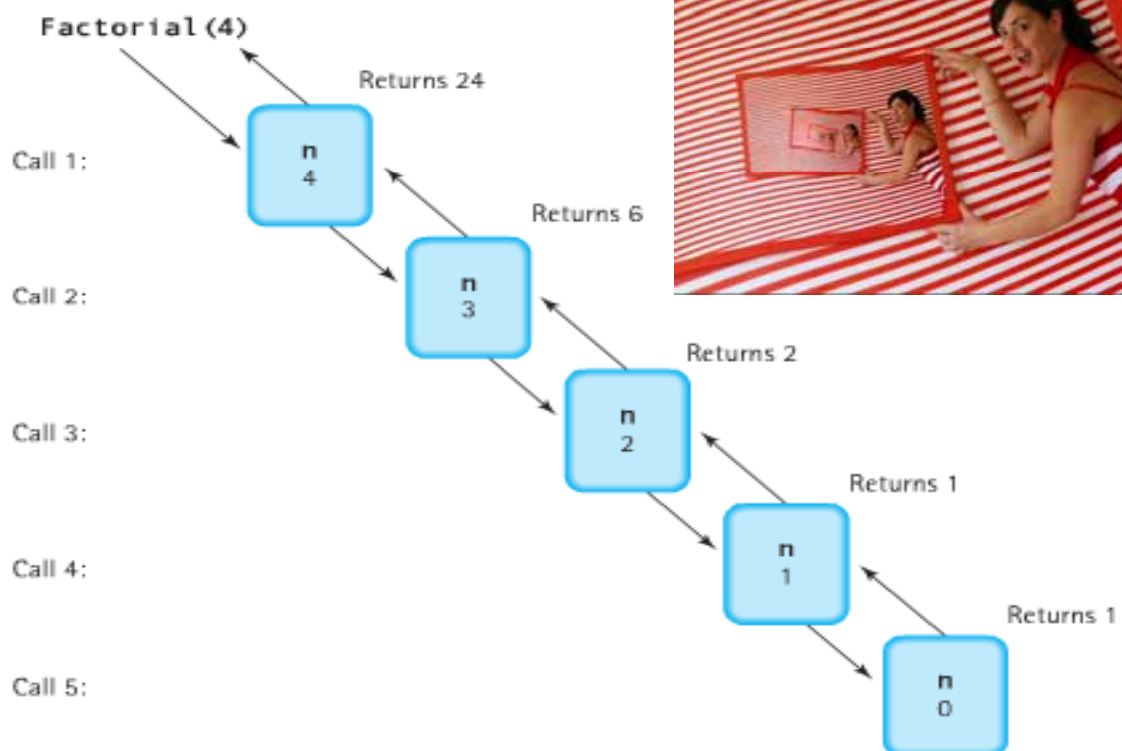
$$v_l(i+1) = e_l(x_{i+1}) \max_k \{v_k(i) a_{kl}\}$$

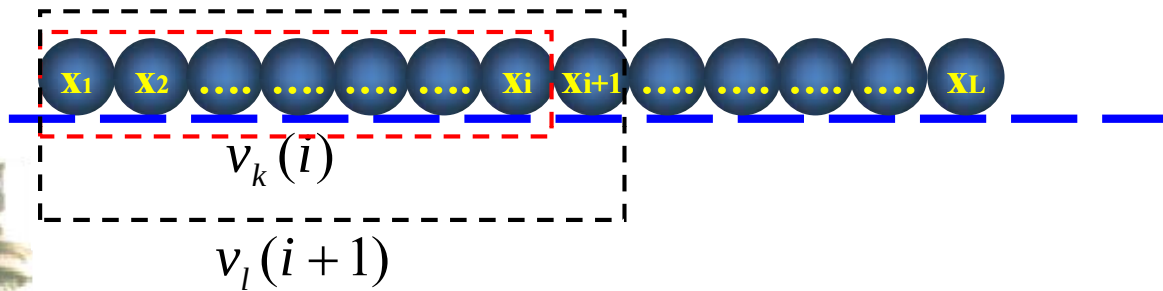
考虑添加起始状态，有：

$$v_0(0) = 1$$



Recursive process





$$v_k(i) = \max_{\{\pi | \pi_i = k\}} P(x_1, x_2, \dots, x_i, \pi)$$

$$\rightarrow v_l(i+1) = e_l(x_{i+1}) \max_k \{v_k(i) a_{kl}\}$$



Viterbi算法的基本步骤

Initialization ($i=0$):

$$v_0(0) = 1$$

$$v_k(0) = 0 \quad \text{for } k > 0$$

Recursion ($i=1, 2, \dots, L$):

$$v_l(i) = e_l(x_i) \max_k \{v_k(i-1) \cdot a_{kl}\}$$

$$ptr_i(l) = \arg \max_k \{v_k(i-1) \cdot a_{kl}\}$$





$$P(x, \pi^*) = \max_k \{v_k(L) \cdot a_{k0}\}$$

$$\pi_L^* = \arg \max_k \{v_k(L) \cdot a_{k0}\}$$

Traceback (i=L, ..., 1):

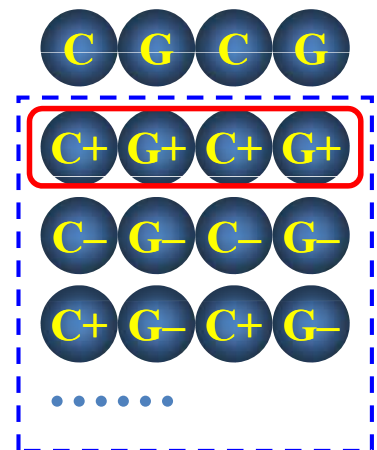
$$\pi_{i-1}^* = ptr_i \left\{ \pi_i^* \right\}$$

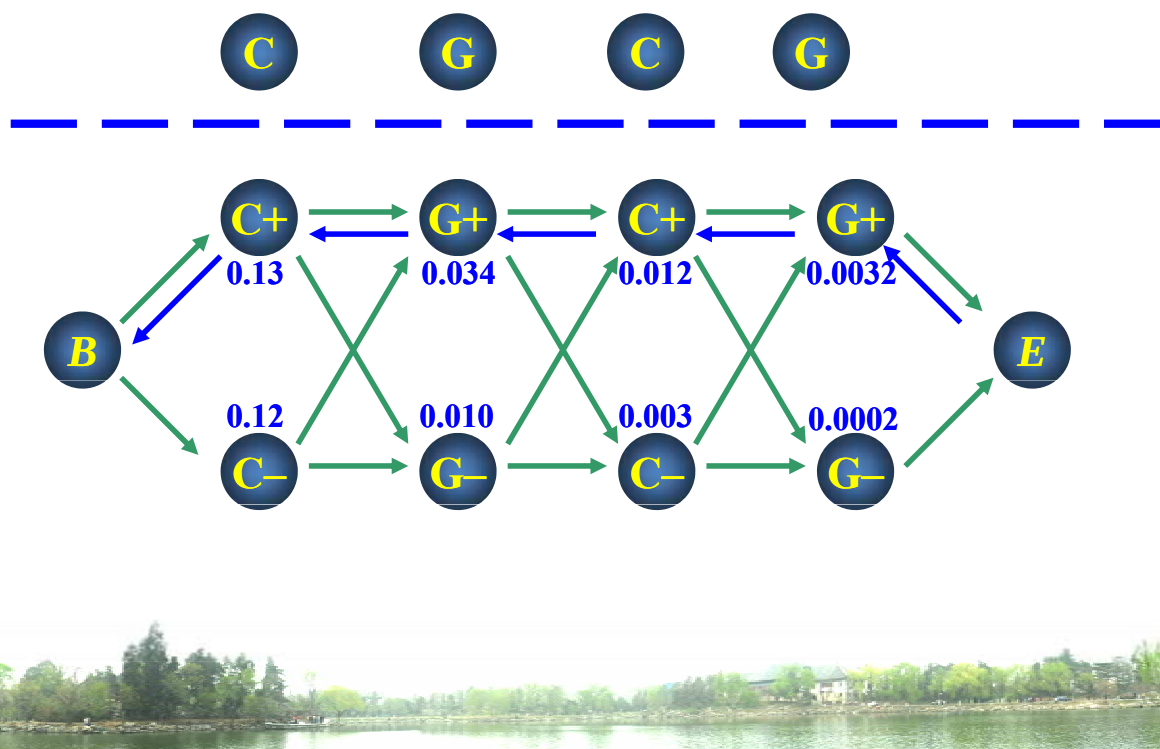
结束



A
C
G
T
A
C
G
T

$v_k(i)$		C	G	C	G
B/E	1	0	0	0	0
A+	0	0	0	0	0
C+	0	0.13	0	0.012	0
G+	0	0	0.034	0	0.0032
T+	0	0	0	0	0
A-	0	0	0	0	0
C-	0	0.12	0	0.0026	0
G-	0	0	0.010	0	0.00021
T-	0	0	0	0	0





With logarithmic scores

目的：为避免乘法运算中小数的影响，可采取对数形式。

Initialization ($i=0$):

$$v_0(0) = 0$$

$$v_k(0) = -\infty \quad \text{for } k > 0$$

Recursion ($i=1, 2, \dots, L$):

$$v_l(i) = \log e_l(x_i) + \max_k \{v_k(i-1) + \log(a_{kl})\}$$

$$ptr_i(l) = \arg \max_k \{v_k(i-1) + \log(a_{kl})\}$$





Termination:

$$P(x, \pi^*) = \max_k \{v_k(L) + \log(a_{k0})\}$$

$$\pi_L^* = \arg \max_k \{v_k(L) + \log(a_{k0})\}$$

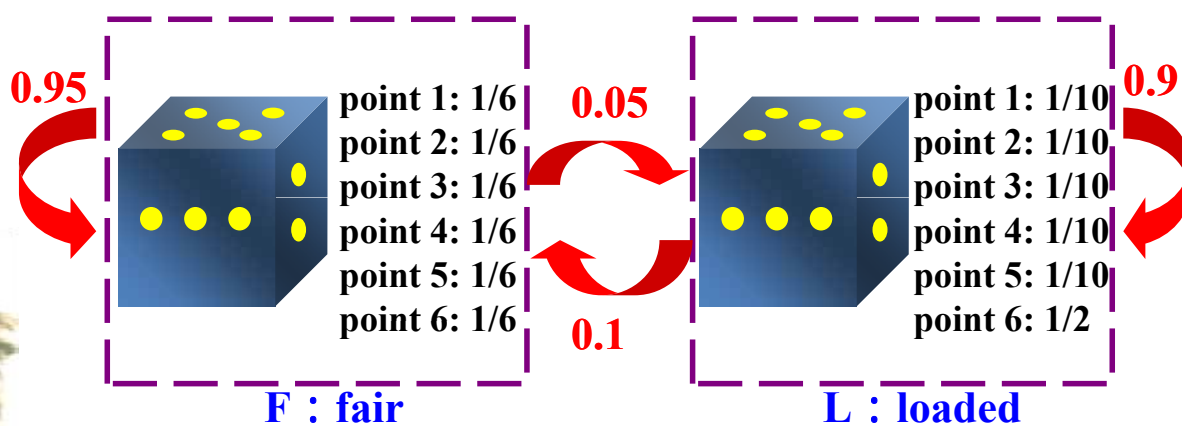
Traceback (i=L, ..., 1):

$$\pi_{i-1}^* = \text{ptr}_i \{\pi_i^*\}$$

结束



Viterbi算法应用举例：赌场骰子模型



状态：{F, L}

输出符号：300次连续投掷的记录

$$a_{FF} = 0.95 \quad a_{FL} = 0.05 \quad a_{LL} = 0.90 \quad a_{LF} = 0.10$$

$$e_F(1) = e_F(2) = e_F(3) = e_F(4) = e_F(5) = e_F(6) = 1/6$$

$$e_L(1) = e_L(2) = e_L(3) = e_L(4) = e_L(5) = 0.1$$

$$e_L(6) = 0.5$$



运用Viterbi算法根据已建立的HMM模型参数，对300次连续投掷的记录进行预测的结果（隐状态序列）之比较：

```
Rolls    315116246446644245311321631164152133625144543631656626566666
Die      FFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFL
Viterbi  FFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFL

Rolls    651166453132651245636664631636663162326455236266666625151631
Die      LLLLLLFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFL
Viterbi  LLLLLLFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFL

Rolls    222555441666566563564324364131513465146353411126414626253356
Die      FFFFFFFFFLLLLLLLLLLLLLLLLLLLLLLLLLLLLLLLLLLLLLLLLL
Viterbi  FFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFL

Rolls    366163666466232534413661661163252562462255265252266435353336
Die      LLLLLLLLLFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFF
Viterbi  LLLLLLLLLLFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFF

Rolls    233121625364414432335163243633665562466662632666612355245242
Die      FFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFL
Viterbi  FFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFFL
```



3. 前向算法(The Forward Algorithm)

计算某DNA片段取某一
Markov过程的可能性

某DNA序列取所有
隐Markov状态过程的可能性

$$P(x) = P(x_1) \prod_{i=2}^L a_{x_{i-1}x_i}$$

识别CpG岛片段与
非CpG岛片段





计算某一DNA序列 $x_1 \dots x_L$ 取全部可能隐状态路径的可能性。考虑所有可能性之和，即序列 $x_1 \dots x_L$ 取所有状态路径的全概率：

$$P(x) = \sum_{\pi} P(x, \pi) = \sum_{\pi} P(\pi | x)$$

上式计算量随状态路径数目呈指数增长，计算量太大。

假定隐状态路径必定是最大可能状态路径，则可以考虑采用最大可能状态路径的 π^* 概率来近似代替全概率 $P(x)$ ：

$$\pi^* = \arg \max_{\pi} P(x, \pi) = \arg \max_{\pi} P(\pi | x)$$



前向算法

实际上，仍可采取类似Viterbi算法的策略，根据动态规划原理，计算每一步的求和，最后得到全概率 $P(x)$ 。

首先定义序列 $x_1 \dots x_L$ 的第 i 个位置上状态 π_i 为 k 的概率(全概率)：

$$f_k(i) = P(x_1 x_2 \dots x_i, \pi_i = k)$$

递归表达式为：

$$f_l(i+1) = e_l(x_{i+1}) \sum_k f_k(i) \cdot a_{kl}$$





前向算法的基本步骤

Initialization ($i=0$):

$$f_0(0) = 1$$

$$f_k(0) = 0 \quad \text{for } k > 0$$

Recursion ($i=1, 2, \dots, L$):

$$f_l(i) = e_l(x_i) \sum_k \{ f_k(i-1) \cdot a_{kl} \}$$

Termination:

$$P(x) = \sum_k f_k(L) \cdot a_{k0}$$

结束



4. 后向算法、后验状态概率

Viterbi算法

前向算法



考察产生观察
序列 $x_1 \dots x_L$ 的概率

- 1、最大概率 $P(x, \pi^*)$
- 2、全概率 $P(x)$

后向算法



考察产生观察
序列 $x_1 \dots x_L$ 中符号
 x_i 的概率

后验状态概率
(Posterior state probability)
 $P(\pi_i = k | x)$



间接求解：

首先考察下列概率

$$\begin{aligned} P(x, \pi_i = k) &= P(x_1 \dots x_L, \pi_i = k) \\ &= P(x_1 \dots x_i, \pi_i = k) \cdot P(x_{i+1} \dots x_L \mid x_1 \dots x_i, \pi_i = k) \\ &= P(x_1 \dots x_i, \pi_i = k) \cdot P(x_{i+1} \dots x_L \mid \pi_i = k) \end{aligned}$$

$$f_k(i)$$

$$b_k(i) = P(x_{i+1} \dots x_L \mid \pi_i = k)$$

$$\Rightarrow P(x_1 \dots x_L, \pi_i = k) = f_k(i) \cdot b_k(i)$$



后向算法(Backward algorithm)的基本步骤

Initialization ($i=L$):

$$b_k(L) = a_{k0} \quad \text{for all } k > 0$$

Recursion ($i=L-1, L-2, \dots, 1$):

$$b_k(i) = \sum_l a_{kl} \cdot e_l(x_{i+1}) \cdot b_l(i+1)$$

Termination:

$$P(x) = \sum_l a_{k0} \cdot e_l(x_1) \cdot b_l(1)$$

$$\text{or } = \sum_l f_l(L) \cdot a_{k0}$$

结束



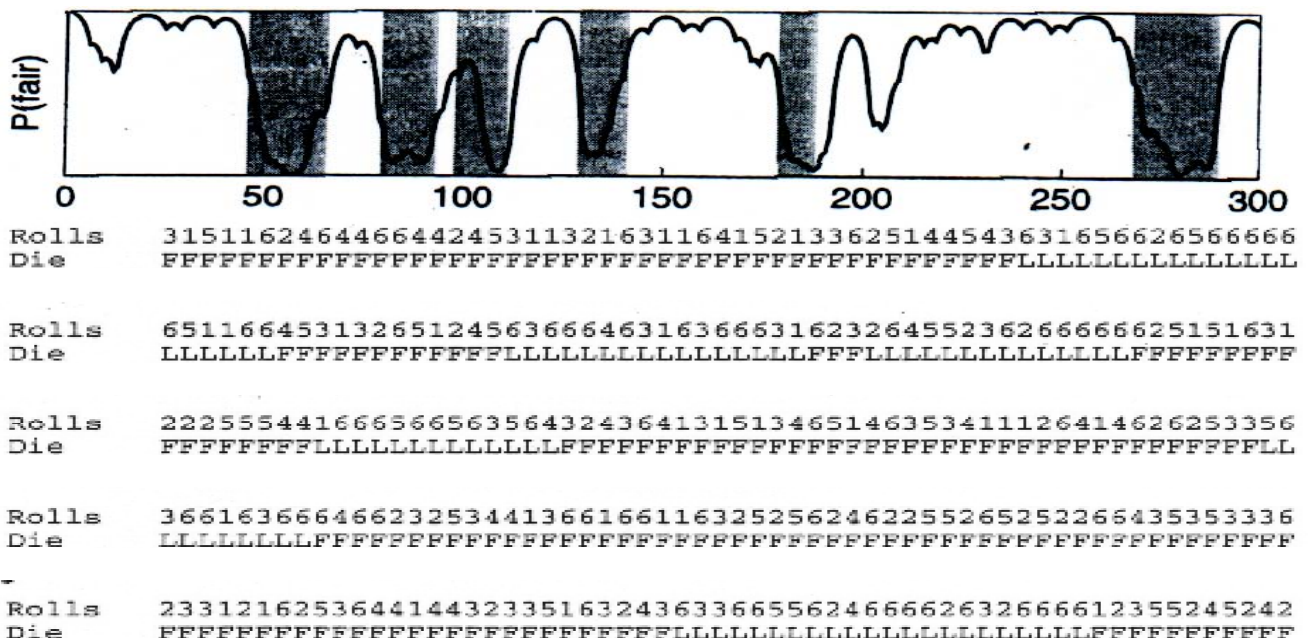
由此可以计算每一步的后验状态概率 (Posterior state probability) :

$$P(\pi_i = k \mid x_1 \dots x_L) = \frac{f_k(i) \cdot b_k(i)}{P(x)}$$



后验状态概率算法举例：赌场骰子模型

假设每一次的观察点数都是由fair die投掷产生的，由此可计算已知观察序列中每一步的后验状态概率：





后验状态概率的应用—— 译码(Decoding)问题

Viterbi算法



多条最大可能的
隐状态路径具有
相近的概率

后向算法

补充方法



由此可以计算由下面每一步后验状态概率决定的序列:

$$\hat{\pi}_i = \arg \max_k P(\pi_i = k | x)$$

用以代替全局最大可能状态路径决定的序列:

$$\pi^* = \arg \max_{\pi} P(x, \pi)$$





§ 7.3 隐Markov模型的 参数估计与结构设计

状态转移概率 a_{kl} :

$$a_{kl} = P(\pi_i = l \mid \pi_{i-1} = k)$$

输出概率 $e_k(b)$:

$$e_k(b) = P(x_i = b \mid \pi_i = k)$$



1. 直接统计得到模型参数 ——训练集状态序列已知的情况

训练数据集 (n 条序列) : $x^1, x^2, \dots, x^n$

当 :

- (1) 训练集数据足够 $\Rightarrow$ 满足最大似然估计的良好条件
- (2) 每一条序列的状态路径已知 (给出所有观察序列的状态的注释信息)

可以直接统计得到HMM模型的各个参数。

例 :

- (1) CpG island : 标注CpG岛区域位点 (起始、终止) 信息的DNA序列
- (2) 蛋白质二级结构 : 标注各类二级结构区域位点信息的氨基酸序列
- (3) 编码基因预测 : 标注基因结构的各个位点



统计训练集中所有状态转移、符号输出的次数（频数）：

$$A_{kl} \quad E_k(b)$$

根据最大似然估计原理，直接计算它们的频率作为概率值：

$$a_{kl} = \frac{A_{kl}}{\sum_{l'} A_{kl'}} \quad e_k(b) = \frac{E_k(b)}{\sum_{b'} E_k(b')}$$

insufficient
training data

Pseudocounts
Modification

HMM的参数

$$a_{kl} \\ e_k(b)$$

2. 自学习得到模型参数 ——训练集状态序列未知的情况

训练数据集（ n 条序列）： $x^1, x^2, \dots, x^n$

当：

- （1）训练集数据足够 $\Rightarrow$ 满足最大似然估计的良好条件
- （2）每一条序列的状态路径未知，无法直接统计得到HMM模型的各个参数。

自学习算法
（优化、迭代）

EM训练算法
（Baum-Welch算法）

Viterbi训练算法



EM训练算法

主要的概率思想

迭代原理：根据当前已知的状态转移概率 a_{kl} 和符号输出概率 $e_k(b)$ ，考虑训练集给定序列的所有可能路径，由此计算期望的状态转移和符号输出的次数 A_{kl} 、 $E_k(b)$ 。再计算新的状态转移概率 a_{kl} 和符号输出概率 $e_k(b)$ ，重复下一步的运算。

迭代终止。



已知：训练数据集（ n 条序列）： $x^1, x^2, \dots, x^n$

假设给定当前已知的状态转移概率 a_{kl} 和符号输出概率 $e_k(b)$

对每一条序列，计算下列概率值：

$$P(\pi_i = k, \pi_{i+1} = l \mid x, \theta) = \frac{f_k(i) \cdot a_{kl} \cdot e_l(x_{i+1}) \cdot b_l(i+1)}{P(x)}$$

其中 θ 表示HMM的模型参数。

由此计算期望的状态转移和符号输出的频数。

$$\Rightarrow A_{kl} = \sum_{j=1}^n \frac{1}{P(x^j)} \sum_i f_k^j(i) \cdot a_{kl} \cdot e_l(x_{i+1}^j) \cdot b_l^j(i+1)$$

$$\Rightarrow E_k(b) = \sum_{j=1}^n \frac{1}{P(x^j)} \sum_{\{i \mid x_i^j = b\}} f_k^j(i) \cdot b_l^j(i)$$



$$a_{kl} = \frac{A_{kl}}{\sum_{l'} A_{kl'}}$$
$$e_k(b) = \frac{E_k(b)}{\sum_{b'} E_k(b')}$$

$$A_{kl} = \sum_{j=1}^n \frac{1}{P(x^j)} \sum_i f_k^j(i) \cdot a_{kl} \cdot e_l(x_{i+1}^j) \cdot b_l^j(i+1)$$

$$E_k(b) = \sum_{j=1}^n \frac{1}{P(x^j)} \sum_{\{i | x_i^j = b\}} f_k^j(i) \cdot b_l^j(i)$$



3. 隐Markov模型的结构设计

具体问题具体分析，需要对问题的深刻理解。

(略)



§ 7.4 隐Markov模型 在计算生物学中的应用

- 两序列比对 (pairwise alignment)
- Profile HMM
 - 分类
 - 数据库搜索
 - 基于profile的多序列比对
- 基因识别
 - Genscan (Burge *et al*, 1997)
 - GeneMark series (Borodovsky *et al*, 1993-2001)
- Motif 搜寻, 信号肽识别, 蛋白质结构预测...



1. 两序列比对

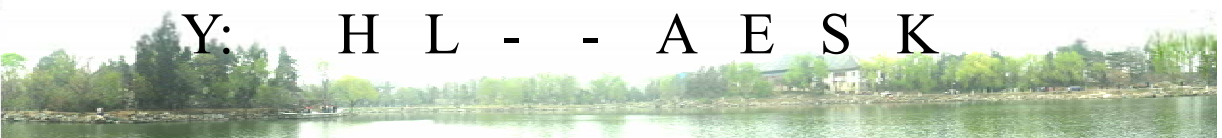
给定两条氨基酸序列X, Y:

X: VLSPADK
Y: HLAESK

假定XY是同源序列。进化过程中，由于插入（Insertion）和缺失（Deletion）的发生，两序列有多种可能的对齐形式，只有一种是真实的比对（alignment）。

一种可能的对齐形式：

X: V L S P A D - K
Y: H L - - A E S K





对alignment进行建模 (a match model) :

X: **V L S P A D - K**

Y: **H L - - A E S K**

M M I_x I_x M M I_y M

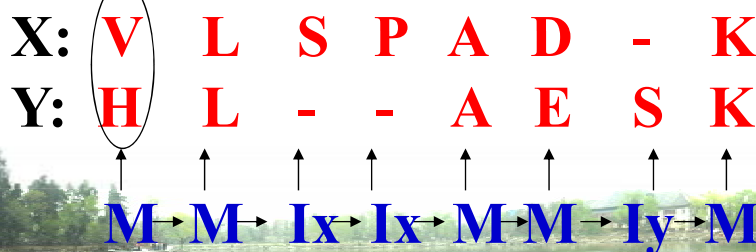
共有三种状态构成了模型的状态空间 {M, I_x, I_y}。

M (match)

I_x (insertion in X)

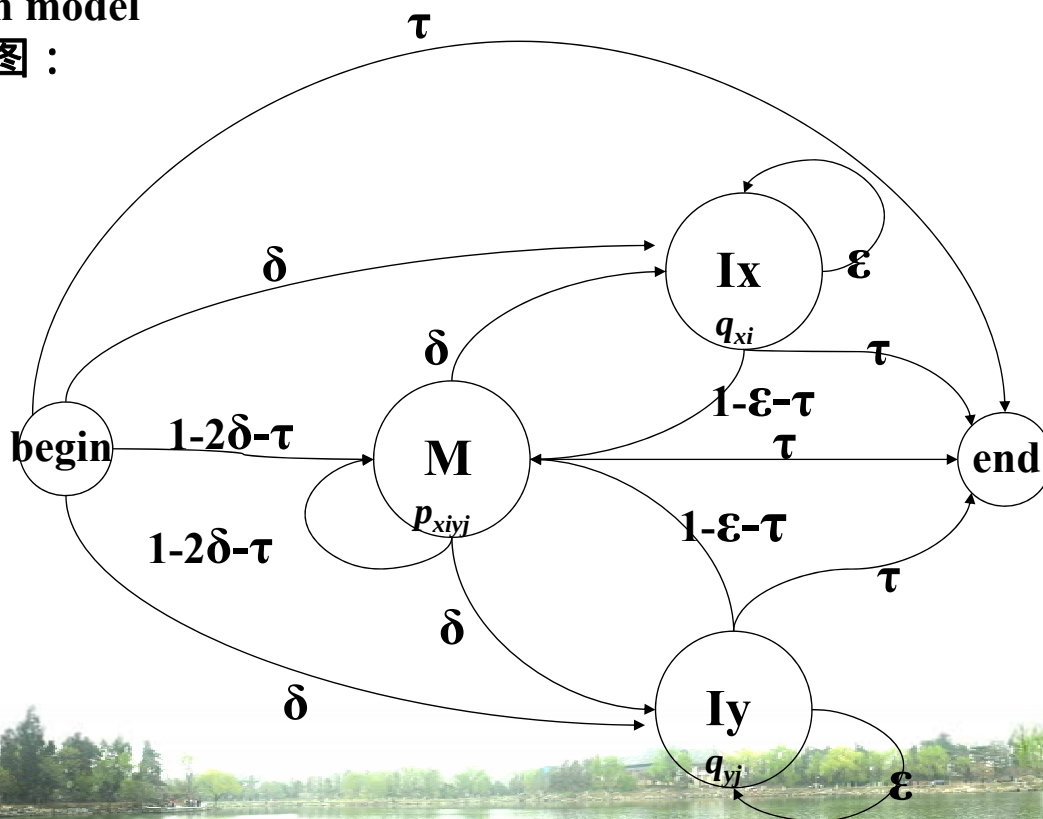
I_y (insertion in Y) ;

任意一个alignment的结果可以看做是对应隐状态链的输出 (emission)。状态之间可以以一定概率转移 (transition)。



Full match model

状态转移图 :





回到给定的 amino acid 序列 X 和 Y :

X: VLSPADK

Y: HLAESK

假定是同源的，如何从众多的对齐形式中找到真实的比对？

所建造的模型是一个概率模型，描述了所有同源序列所有可能的对齐形式，并对每一个赋予概率：

$$P(x \langle \rangle y, \pi)$$

选择概率最大的为最终比对：

$$\pi^* = \arg \max_{\pi} P(x \langle \rangle y, \pi) = \arg \max_{\pi} \underbrace{P(x \langle \rangle y | \pi)}_{\text{emissions}} \cdot \underbrace{P(\pi)}_{\text{transitions}}$$



动态规划算法寻找最优路径 (Viterbi algorithm).

注意: 由于是两条序列，搜索空间为二维 (i,j), i=0,...,7;
j=0,...,6.

递归式：

$$v^M(i, j) = p_{x_i y_j} \max \begin{cases} (1 - 2\delta - t)v^M(i-1, j-1), \\ (1 - \varepsilon - t)v^{Ix}(i-1, j-1), \\ (1 - \varepsilon - t)v^{Iy}(i-1, j-1); \end{cases}$$

$$v^{Ix}(i, j) = q_{x_i} \max \begin{cases} \delta v^M(i-1, j), \\ \varepsilon v^{Ix}(i-1, j); \end{cases}$$

$$v^{Iy}(i, j) = q_{y_j} \max \begin{cases} \delta v^M(i, j-1), \\ \varepsilon v^{Iy}(i, j-1); \end{cases}$$



再回到序列X, Y:

X: VLSPADK

Y: HLAESK

问题：它们是否同源；或者它们是同源的可能性是多少？

单一的match 模型给出了假定是同源的情况下，X,Y最可能具有的比对形式 π^* 。以这一比对形式出现的同源序列的概率为：

$$P(x \triangleleft y, \pi^* | M)$$

现在考察这一比对的significance，需要另外一个模型——随机模型R。R刻画了所有可能的两序列对，并赋予一定的概率

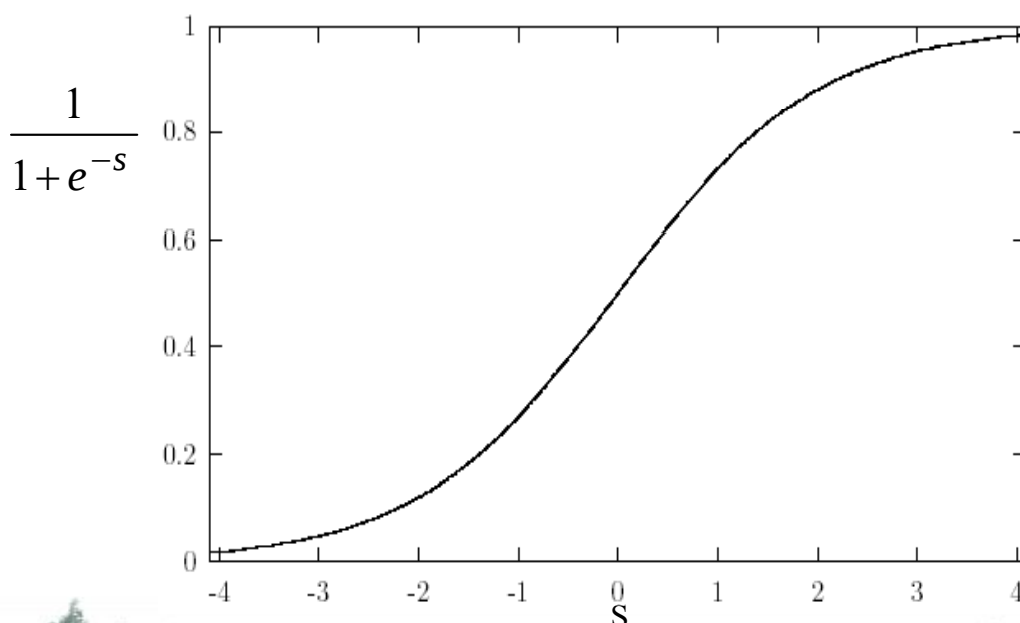
$$P(x, y | R)$$

我们考察两概率的log-odd 值：

$$S = \log \frac{P(x \triangleleft y, \pi^* | M)}{P(x, y | R)} + S'$$



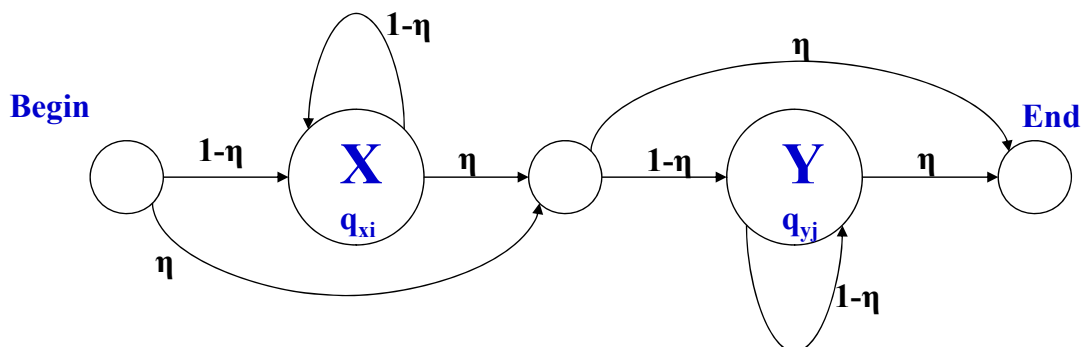
Logistic 函数





Full random model

状态转移图:



重新做比对，考虑Random model的影响，即

$$\pi^* = \arg \max_{\pi} \log \left(\frac{P(x \text{ <> } y, \pi)}{P(x, y | R)} \right) = \arg \max_{\pi} \log \left(\frac{P(x \text{ <> } y | \pi) \cdot P(\pi)}{P(x, y | R)} \right)$$

根据模型假设：

$$P(x, y | R) = \eta(1-\eta)^7 \prod_{i=1}^7 q_{x_i} \cdot \eta(1-\eta)^6 \prod_{j=1}^6 q_{y_j}$$

将其各项融进match model 的对应项中：

$$\mathbf{M}: \log \frac{p_{x_i y_j}}{q_{x_i} q_{y_j}} + \log \frac{a_{IM}}{(1-\eta)^2}$$

$$\text{定义: } s(a, b) = \log \frac{p_{ab}}{q_a q_b} + \log \frac{(1-2\delta-\tau)}{(1-\eta)^2}$$

$$\mathbf{Ix}: \log \frac{a_{Ix}}{1-\eta}$$

$$d = -\log \frac{\delta(1-\varepsilon-\tau)}{(1-\eta)(1-2\delta-\tau)}$$

$$\mathbf{Iy}: \log \frac{a_{Iy}}{1-\eta}$$

$$e = -\log \frac{\varepsilon}{1-\tau}$$



比对的分值：

X: V L S P A D - K

Y: H L - - A E S K

$$\begin{aligned} & S(V,H)+S(L,L)-d-(2-1)e+S(A,A)+S(D,E)-d+S(K,K) \\ & =-4+5-10-2+5+2-10+6 \\ & =-8 \end{aligned}$$



In fact, the search and alignment applications constitute probably the best-known use of HMM for biological sequence analysis. However, HMM theory should be emphasized with its much broader applicability, which goes far beyond that of sequence alignment.

—— 《Biological sequence analysis》



2. 蛋白质家族的统计描述 (Profile HMM)

一个蛋白家族 (family) 包含有众多的成员蛋白(member)。成员蛋白之间序列相差较大，但结构比较保守，功能也相似或相关，因为它们来自共同的祖先。

蛋白家族的profile即指，对所有成员蛋白共同具有的特征进行统计性的描述。在寻找新的成员蛋白时，用这个profile对数据库中的每一个candidate进行打分，以此判断其是否为该家族成员。

用HMM的方法对某一蛋白家族进行统计描述，称为profile HMM.



CSRC_HUMAN	1	42	KLGGQCFGEVVMGTW-----NGTTRVAIKTLKP-----GTMSPFAF---LQEAQV
CABL_HUMAN	1	43	KLGGQYGEVYEGVW-----KKYSLTVAVKTLKE-----DTMEVEEF---LKEAAV
EPH_HUMAN	1	47	-IGEGFGEVYRGTLR---LPSQDCKTVAIKTLKDT---SPGGQWWNF---LREATI
FER_HUMAN	1	44	LLGKGNFGEVYKGTL-----KDKTSAVKTCKED---LPQELKIKF---LQEAKI
IR_HUMAN	1	48	--GQGSFGMVYEGNARDI--IKGEAETRVAVKTVNES---ASLRERIEF---LNEASV
CROS_HUMAN	1	49	--GSGAFGEVYEGTAVDI-LGVGSGEIKVAVKTLKKG---STDQEKIEF---LKEAHL
TRK_HUMAN	1	47	--GEGAFGKVF LAECHNL--LPEQDKMLVAVKALKE---ASESARQDF---QREAEI
BFGFR_HUMAN	1	50	--GEGCFGQVLAFAIGLDKDKPNRVTKVAVKMLKSD---ATEKDLSDL---ISEMEM
RET_HUMAN	1	48	--GEGFEGKVVKATAFHL--KGRAGYTTVAVKMLKEN---ASPSELRLD---LSEFNV
EGFR_HUMAN	1	47	--GSGAGTVYKGLWIP---EGEKVKIPVAIKELREA---TSPKANKEI---LDEAYV
SCH9_YEAST	1	47	LLGKCTFGQVYQVKK-----KDTQRIYAMKVL SKKVI--VKKNEIAHT---TGERNI
S6K_70K_RAT	1	50	-LGKGGYGVVFQVRKV---TGANTGKIPAMKVLKKAMI--VRNAKDTAHT---KAERNI
PKC1_YEAST	1	47	VLGKGNFGKVLKSKS-----KNTDRLCAIKVLKKNI--IQNHDIESA---RAEKKV
CGMPK_HUMAN	1	47	-LGVGCFGRVELVQL-----KSEESKTFAMKILKKRHI--VDTRQQEHI---RSEKQI
PRK_RABBIT	1	53	ILGRGVSSVVRRCIH-----KPTCKEYAVKIIDVTGGGSFSAEEVQELREATLKEVDI
CSRC_HUMAN	43	80	M---KKL-R-HEKLVQLYAVVSE-EPIYIVTEYMSKGSLLDPLK
CABL_HUMAN	44	80	M---KEI-K-HPNLVQLLGVCTREPPFYIITEFMTYGNLLDY--
EPH_HUMAN	48	80	M---GQF-S-HPHILHLEGVVTKRKPIMIITEFMENGA-----
FER_HUMAN	45	80	L---KQY-D-HPNIVKLIGVCTQRQPVYIIMELVSGGDFLT---
IR_HUMAN	49	80	M---KGF-T-CHHVVRLLGVVSKGQPTLVVMEI MAHG-----
CROS_HUMAN	50	80	M---SKF-N-HPNIIKQLGVCLLNEPQYIILELMEG-----
TRK_HUMAN	48	80	L---TML-Q-HQHIVRFVGCTEGRPLL MVFEYMRHGD-----
BFGFR_HUMAN	51	80	M---KMIG-KHNIIINLLGACTQDGP LYVIVEYA-----
RET_HUMAN	49	80	L---KQV-N-HPHVIKLYGACSQDGPLLLI V EYAKYG-----
EGFR_HUMAN	48	80	M---ASV-D-NPHVCRLIGICLT-STVQLITQLMPFGCL-----
SCH9_YEAST	48	80	LVTTASK-S-SPFIVGLKFSFQTPDLYLVTDYMS-----
S6K_70K_RAT	51	80	L---EEV-K-HPFIVDLIYAFQTGCKLYLILEYLS-----
PKC1_YEAST	48	80	FLLATKT-K-HPFLTNLYCSFQTE NRIYFAMEFIG-----
CGMPK_HUMAN	48	80	M---QGA-H-SDFIVRLYRTFKDSKYLYMLMEACL GGE-----
PRK_RABBIT	54	80	L---RKV-SCHPNIIQLKDYETNTFFFLVF-----

An alignment of 15 kinase: human, rabbit, rat & yeast; retrieved from PAPIA server.



怎样建profile HMM?

首先对已知的成员蛋白做多序列比对 (基于结构或基于序列).

注意：必须保证得到的比对是正确的，因为下一步将对其进行建模，然后搜寻新的成员蛋白。

然后对比对好的多重序列进行建模 (consensus modeling)。对于序列的每一个位置都有三种可能的状态：match、insertion、deletion。

Match状态（或者说main状态）处于哪些位置必须确定下来，以此确定模型的拓扑结构。

比如：

```
WMGTW-----NGTTRVA
YEGVW-----KKYSLTVA
YRGTLR----LPSQDCKTVA
YKGT-----KDKTSVA
YEGNARDI--IKGEAETRVA
YEGTAVDI-LGVGSGEIKVA
FLAECHNL--LPEQDKMLVA
VLAEAIGLDKDKPNRVTKVA
VKATAFHL--KGRAGYTTVA
YKGLWIP--EGEKVKIPVA
YQVKK-----KDTQRIYA
FQVRKV----TGANTGKIFA
ILSKS-----KNTDRLCA
ELVQL-----KSEESKTFA
RRCIH-----KPTCKEYA
*****
```

profile HMM 的状态转移图

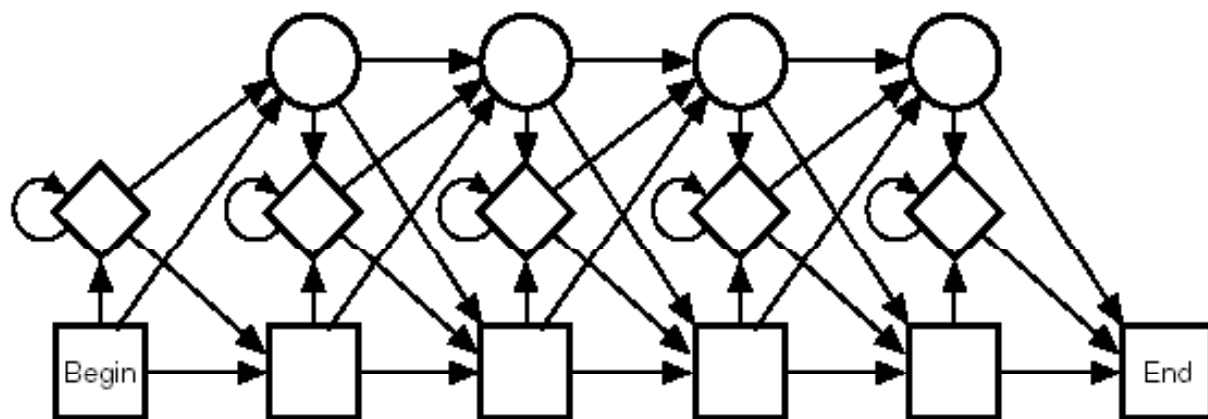


Figure from *Computational methods in molecular biology*, Chap. 4, by Anders Krogh.



: Match



: Insertion



: Deletion



参数估计

(Transition probabilities and emission probabilities)

直接从多道比对中，通过统计得出。

由于样本有限，一些允许的事件出现的次数为零——运用 Laplace 准则: 对每一个可能的事件都事先加一 (pseudo-count=1)。

例如：

```
WMGTW-----NGTTRVA
YEGVW-----KKYSLTVA
YRGTLR----LPSQDCKTVA
YKGTL-----KDKTSVA
YEGNARDI--IKGEAETRVA
YEGTAVDI-LGVGSGEIKVA
FLAECHNL--LPEQDKMLVA
VLAEAIGLDKDKPNRVTKVA
VKATAFHL--KGRAGYTTVA
YKGLWIP--EGEKVKIPVA
YQVKK-----KDTQRIYA
FQVRKV---TGANTGKIFA
ILSKS-----KNTDRLCA
ELVQL-----KSEESKTFA
RRCIH-----KPTCKEYA
*****
```

position 1 2 3 4 5...



$$e_{M_1}(W) = 2 / 35; \quad a_{M_1M_2} = 16 / 18;$$

$$e_{M_1}(Y) = 8 / 35; \quad a_{M_1I_1} = 1 / 18;$$

$$e_{M_1}(F) = 3 / 35; \quad a_{M_1D_2} = 1 / 18.$$

$$e_{M_1}(V) = 3 / 35;$$

$$e_{M_1}(I) = 2 / 35;$$

$$e_{M_1}(E) = 2 / 35;$$

$$e_{M_1}(R) = 2 / 35;$$

$$\text{其它的为 } 1 / 35. \quad \dots$$



数据库搜寻

寻找新的成员蛋白(significant matches to the profile):

首先运用动态规划方法将每一个可能的匹配 (match) 同已有的比对 (profile) 对齐。

$$\pi^* = \arg \max_{\pi} P(x, \pi \mid \text{profile})$$

这个匹配 是家族成员的似然性(likelihood)即是：

$$P(x, \pi^* \mid \text{profile})$$

为了考察匹配的significance，需要一随机模型进行比较——log-odd 版本的Viterbi算法。

一个标准的随机模型产生这条序列 (match) 的概率为：

$$P(x \mid R) = \prod_i q_{x_i}$$



Log-odd 版本的Viterbi 递归式：

$$V_j^M(i) = \log \frac{e_{M_j}(x_i)}{q_{x_i}} + \max \begin{cases} V_{j-1}^M(i-1) + \log a_{M_{j-1}M_j}, \\ V_{j-1}^I(i-1) + \log a_{I_{j-1}M_j}, \\ V_{j-1}^D(i-1) + \log a_{D_{j-1}M_j}; \end{cases}$$

$$V_j^I(i) = \log \frac{e_{I_j}(x_i)}{q_{x_i}} + \max \begin{cases} V_j^M(i-1) + \log a_{M_jI_j}, \\ V_j^I(i-1) + \log a_{I_jI_j}, \\ V_j^D(i-1) + \log a_{D_jI_j}; \end{cases}$$

$$V_j^D(i) = \max \begin{cases} V_{j-1}^M(i) + \log a_{M_{j-1}D_j}, \\ V_{j-1}^I(i) + \log a_{I_{j-1}D_j}, \\ V_{j-1}^D(i) + \log a_{D_{j-1}D_j}; \end{cases}$$

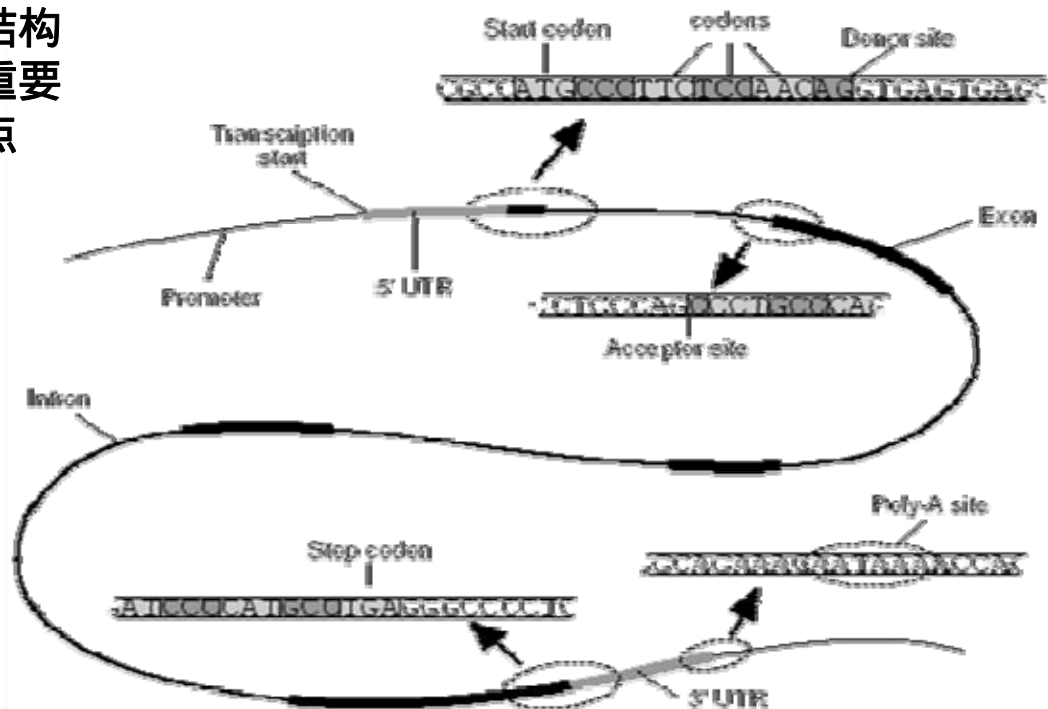


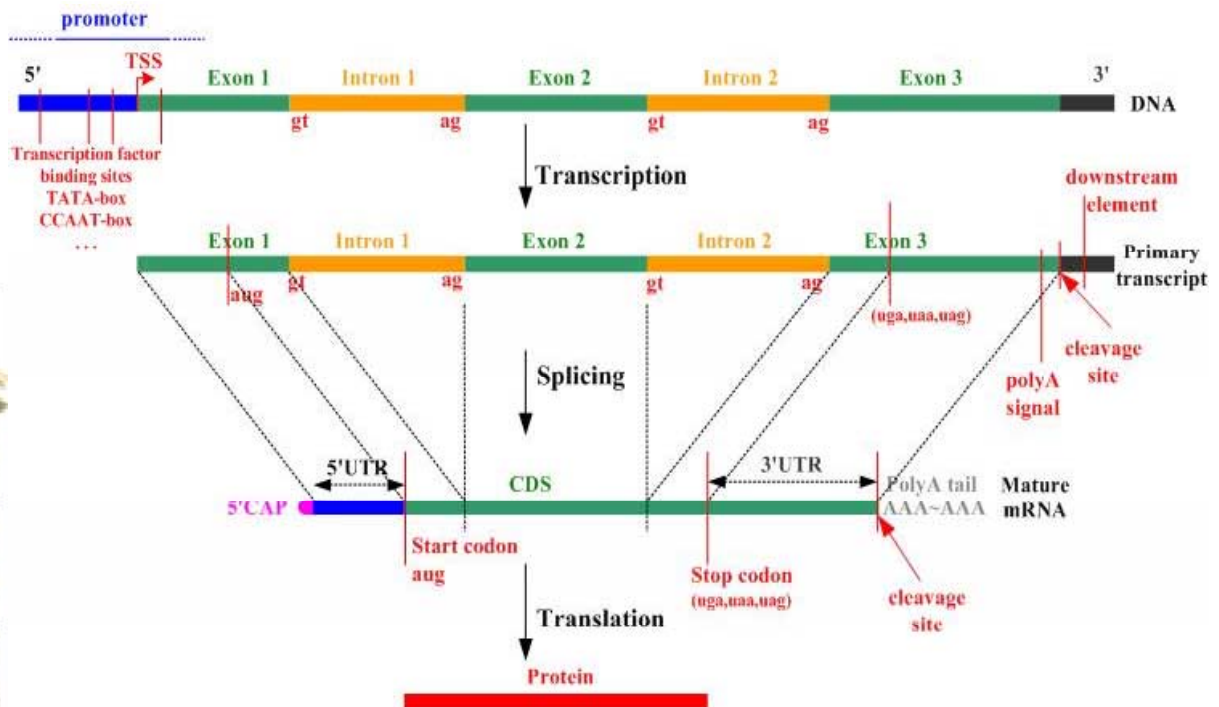
3. 基于HMM的基因识别与预测

- GENSCAN (Burge 1997)
- FGENESH (Solovyev 1997)
- HMMgene (Krogh 1997)
- GENIE (Kulp 1996)
- GENMARK (Borodovsky & McIninch 1993)
- VEIL (Henderson, Salzberg, & Fasman 1997)



真核基因结构
以及一些重要
的调控位点





<http://online.itp.ucsb.edu/online/infobio01/burge/>

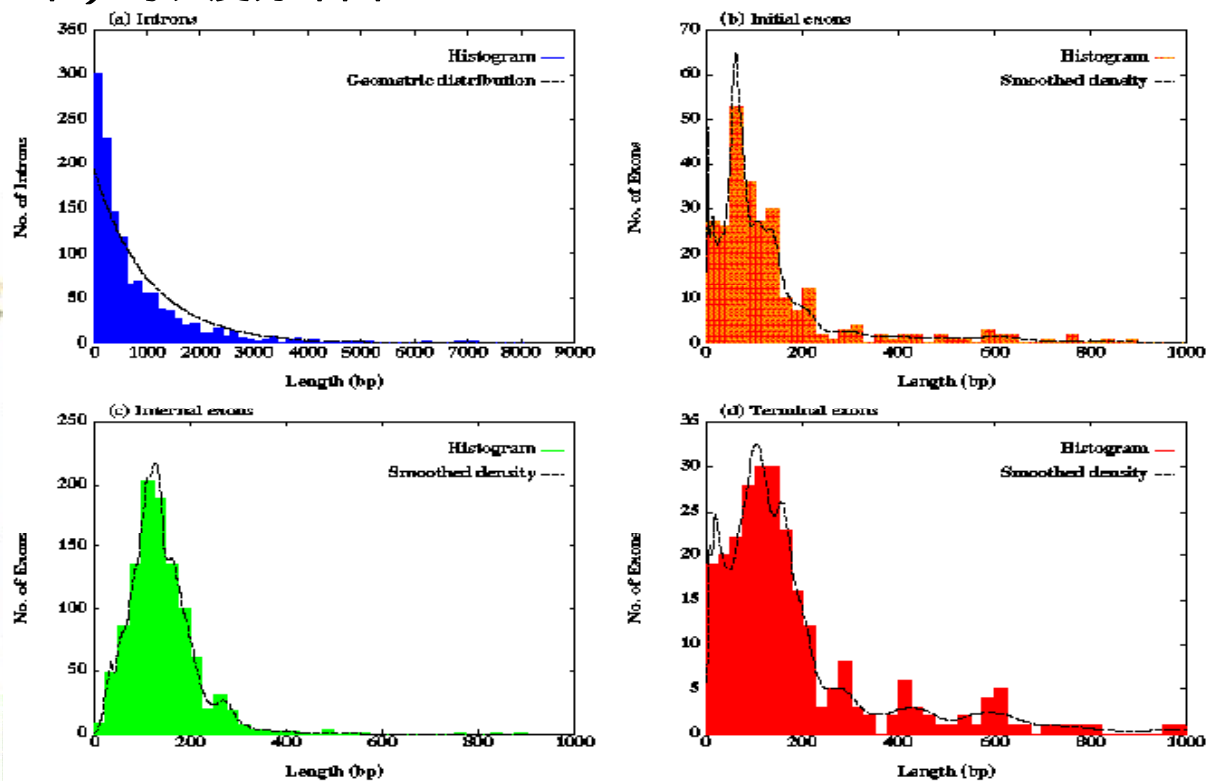


初步统计

- 脊椎动物基因平均长度为 30kb。
- 编码区大约长 1kb。
- 外显子长度变化很大，从几十到几千个碱基。
- 5'非翻译区 (5' UTR) 大约长 750 bp。
- 3'非翻译区 (3' UTR) 大约长 450 bp。



人基因的内含子和三类外显子（第一个、中间的、以及最后一个）的长度分布图



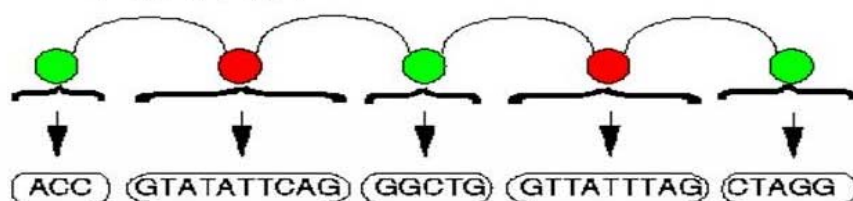
半马氏以及半隐马氏

半马氏链



状态链为Markov链；
每个状态有特定的停留时间。

半隐马氏链



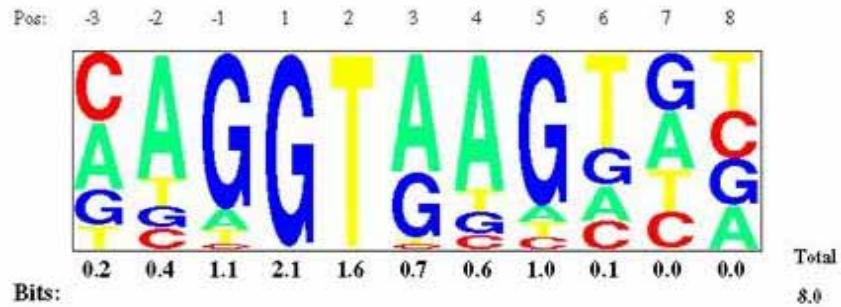
状态不能被观察到；
观察序列由隐状态产生，长度同隐
状态对应的停留时间。



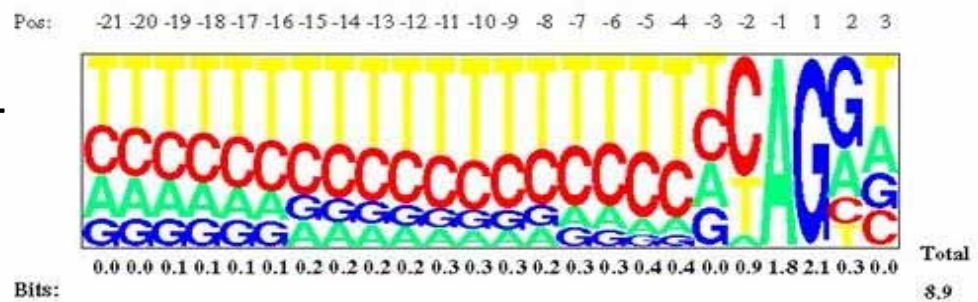
人类基因剪接位点信号



5' 剪接信号

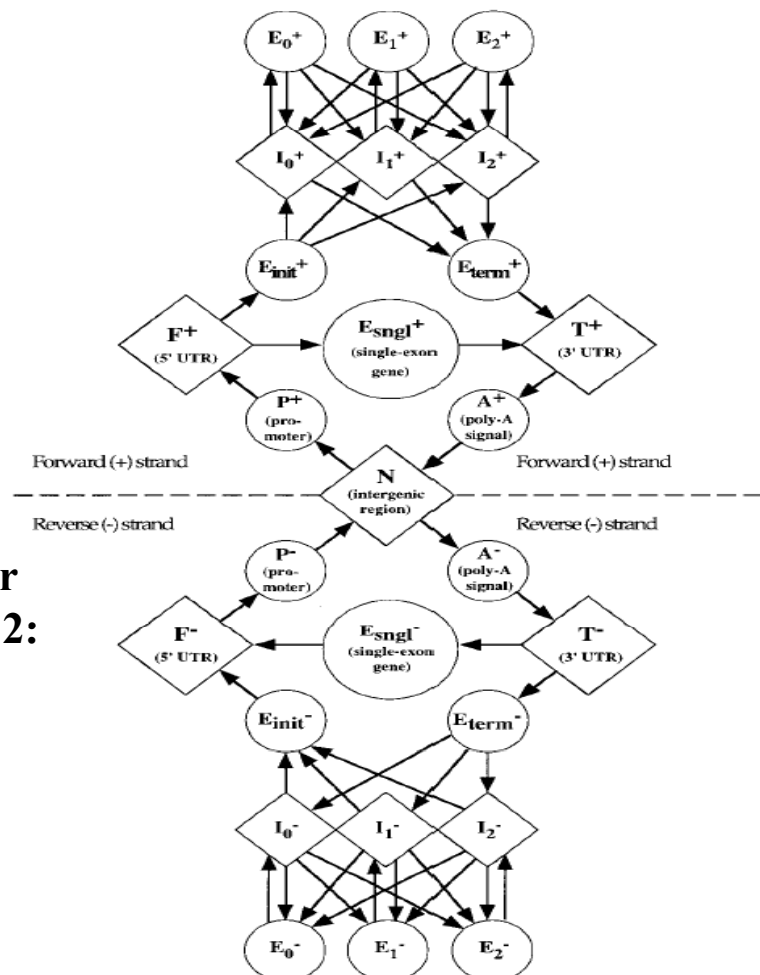


3' 剪接信号



- N – 基因间区
 - P – 启动子
 - F – 5' 非翻译区
 - E_{sngl} – 单外显子基因
 - E_{init} – 第一个外显子
 - E_k – k相位的中间外显子
 - E_{term} – 最后一个外显子
 - I_k – k相位的内含子
- (0: between codons; 1: after the first base of a codon; 2: after the second base of a codon)

Burge, C. & Karlin, S, J.
Mol. Biol. (1997) 268, 78-94.



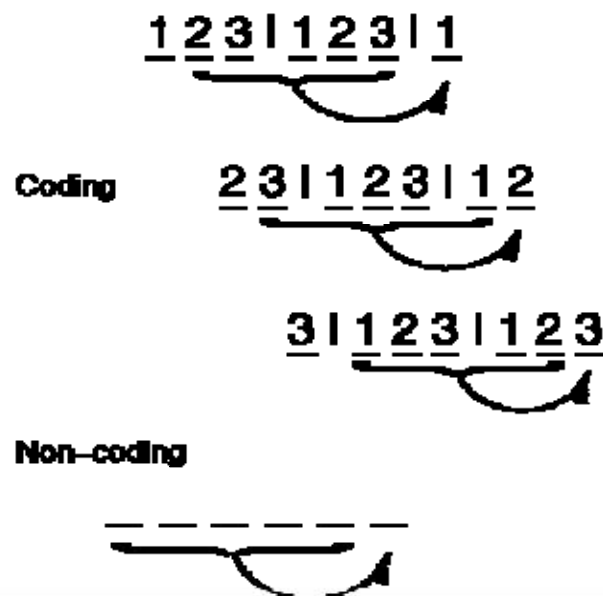


GenScan 信号建模

- WMM (Weight Matrix 模型)
 - Poly A;
 - 翻译起始和终止位点;
 - 启动子;
- WAM (Weight Array 模型)
acceptor 剪接位点 (2阶)
- MDD (Maximal Dependence Decomposition)
donor 剪切位点



编码与非编码区的建模





GenScan 的特点

- 1、对两条链（正，负）同时建模；
- 2、每一个状态可以产生一串的核苷酸序列；
- 3、考虑内含子与外显子的长度分布；
- 4、对剪切位点做单独的建模；
- 5、直接从带有标注（annotation）的学习集中提取参数。



4. 隐Markov模型的讨论

HMM的局限

自由参数过多、对学习集要求较高

模型的物理意义？

HMM的优势

坚实的统计学基础

有效的学习算法