



生物信息学

第八章 序列模式识别 (1)

本章内容提要



- 序列模式
- 位点特异性打分矩阵/权重矩阵模型
 - ✿ Position Specific Scoring Matrix (PSSM),
Weight Matrix Model (WMM)
- 模体发现: Gibbs Sampler等
- 马尔科夫及隐马尔科夫模型
- 翻译后修饰位点预测

序列模式



- ❑ 功能结构域, **functional domain**
- ❑ 模体, **motif**
- ❑ 模块, **BLOCK**
- ❑ 模式, **pattern/profile**

功能结构域/Domain



- ❑ 具有完整的、独立的三级结构
- ❑ 具有特定的生物学功能
- ❑ 一般长度，几十到几百个氨基酸
- ❑ 允许插入/缺失，即允许存在gap



```
KDC2_DROME/1-5 KHLNWNDVYS----KKLKPPILPDVH-----HDGDTKNFD-DYPEKDWKP-----
CONSENSUS/80% ttlsWptl.....ph.sPhhP.lp.....s..Dhp.FD.thspt.....
CONSENSUS/65% psIsWcpl.p....pplpPPahPplp.....s.tDsssFD.casppts.....
CONSENSUS/50% +sIDW-clpp....+clpPPFpPplp.....uspDsoNFD.-FTcpsss.....

KDC2_DROME/1-5 -AKAVDQ-----RDLYYFND
CONSENSUS/80% ..p.....a.shsh..
CONSENSUS/65% .hssssp...t.....Ftuasaht
CONSENSUS/50% .hssssshhpshppp.....FtGFoass
```

模体/Motif



- ❑ 不具有独立的三级结构
- ❑ 具有特定的生物学功能：结合，修饰，细胞亚定位，维持结构，等
- ❑ 长度一般几个到几十个氨基酸或者碱基；
- ❑ 例如，SUMO化的序列模体： Ψ -K-X-E (Ψ :A, I, L, V, M, F, P; X: 任意氨基酸)

Functional site class:

Sumoylation site

Functional site description:

Motif recognised for modification by SUMO-1

ELM(s):

MOD_SUMO

MOD_SUMO description:

Motif recognised for modification by SUMO-1

Pattern:

[VILMAFP]K.E

模块/BLOCK



- ❑ 几个到几十个氨基酸
- ❑ 无gap，从全局多序列比对的结果直接处理得到
- ❑ 描述蛋白质家族或者一类蛋白质的序列保守性



模式/Pattern/Profile



- ❑ 在算法上用来描述一类功能结构域，模体或者模块的表示方式
- ❑ 根据序列数据，构建的预测模型
- ❑ 数据形式：概率表示
- ❑ 用来预测新的可能符合特定模式的序列
- ❑ 例如，直接将 Ψ -K-X-E视为SUMO化位点的，普适的“模式”，则可以预测所有包含该模式的蛋白质序列

位点特异性打分矩阵



- ❑ **Position Specific Scoring Matrix (PSSM)/ Weight Matrix Model (WMM)**
- ❑ **对蛋白质家族进行多序列比对分析，发现结果中保守的BLOCK**
- ❑ **根据BLOCK序列推导相应的PSSM**
- ❑ **不考虑gap的影响**
- ❑ **BLOCK长度一般在几个~几十个残基/碱基**



PSSM Viewer

Amino Acid Explorer

PSSM Viewer Help

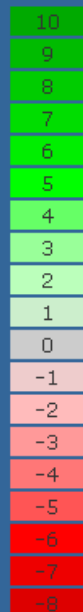
CDD Help

Show Color Key

Course Main Page

Questions or comments

Scores



cd02249 : Zinc finger, ZZ type. Zinc finger present in dystrophin, ...

table those positions where the is

table this feature:

Click on any score to compare the two residues.

Click on any column to sort the matrix by that column's scores.

P - consensus sequence position C - consensus sequence residue

Master:			1TOT_A										View Master in CD													
P	C	Master	A	G	I	L	V	M	F	W	P	C	S	T	Y	N	Q	H	K	R	D	E				
1	Y	7 - Y	-3	-5	2	-2	1	1	3	-2	0	-4	0	-3	6	0	-3	-2	-3	1	-4	-4				
2	S	8 - T	0	-4	2	-1	0	-2	1	-4	0	1	3	3	-3	-3	-2	-3	0	1	-3	-1				
3	C	9 - C	-2	-5	-3	-3	-3	-3	-4	-4	-5	11	-3	-3	-4	-5	-5	-5	-5	-6	-6	-6				
4	D	10 - N	-1	-3	-5	-5	-5	-4	-5	-6	-4	-5	-2	-1	-4	4	-2	3	0	-1	6	-1				
5	G	11 - E	0	4	1	-2	-1	-3	-4	-4	-4	4	-1	0	0	1	-3	3	-3	-4	-3	0				
6	C	12 - C	-2	-5	-3	-3	-3	-3	-4	-4	-5	11	-3	-3	-4	-5	-5	-5	-5	-6	-6	-6				
7	L	13 - K	-3	3	-1	2	-1	-2	-4	-4	-4	-4	0	-3	-4	1	2	-3	1	0	-1	-1				
8	K	14 - H	-1	1	0	-1	-1	2	-3	-4	-4	-4	1	1	0	1	1	1	2	-2	-3	0				
9	P	Gap	-1	0	-3	-1	-2	-3	-3	-4	5	-4	-2	0	1	1	-2	2	-2	-3	1	3				
10	I	15 - H	-3	-5	6	1	1	-1	2	-4	-5	-3	-1	-3	-2	-1	-4	2	-4	-4	-5	-4				
P	C	Master	A	G	I	L	V	M	F	W	P	C	S	T	Y	N	Q	H	K	R	D	E				
11	V	16 - V	0	-4	2	-1	1	3	-2	5	1	-3	0	1	2	-1	0	-3	-3	-1	-1	-3				
12	G	17 - E	-2	5	-5	-5	-4	-4	-5	-4	2	-4	-2	-3	-4	1	-2	-3	-2	-1	-2	3				
13	V	18 - T	-1	-4	0	-1	2	0	0	-4	3	-4	0	0	0	0	-2	1	1	2	-3	1				
14	R	19 - R	-3															0	7	-4	-2					
15	Y	20 - W	-4															-4	-4	-6	-5					
16	H	21 - H	-1															4	2	-2	1					
17	C	22 - C	-2															-5	-5	-5	-5					
18	L	23 - T	1															1	1	-1	0					
19	V	24 - V	-2	-4	0	-3	4	1	-4	-5	-4	-4	0	2	-4	1	0	-3	-1	1	2	1				
20	C	25 - C	0	-4	-3	-3	-3	-3	-4	-4	-5	11	-3	-3	-4	-4	-4	-5	-1	-5	-5	-5				
P	C	Master	A	G	I	L	V	M	F	W	P	C	S	T	Y	N	O	H	K	R	D	E				

锌指功能结构域的PSSM

锌指功能结构域的PSSM

BLOCK -> PSSM



二十种
氨基酸

1 2 3 4 5 11
 ... G D S F H Q F V S H G ...
 ... S D A F H Q Y I S F G ...
 ... G D S Y W N F L S F G ...
 ... S D S F H Q F M S F G ...
 ... G D S Y W N Y A S F G ...

代表每一列

	A	C	D	E	F	G	H	I	K...	S	T...
1						Log(3)				Log(2)	
2			Log(5)								
3											
4											
5											

矩阵中的数值：当前位置上，某种氨基酸出现的频率的log值

第二种PSSM



- 每一个位置上显示每种氨基酸或者碱基出现的频率

四种碱基

碱基的位置

>ABF17121SCPD

A	0	0	11	2	3	7	3	5	5	14	0	1
T	14	0	2	1	6	2	7	3	0	0	0	0
G	0	0	0	0	2	3	1	2	3	0	1	13
C	0	14	1	11	3	2	3	4	6	0	13	0.

第三种PSSM



- 每一个位置显示氨基酸/碱基出现的概率



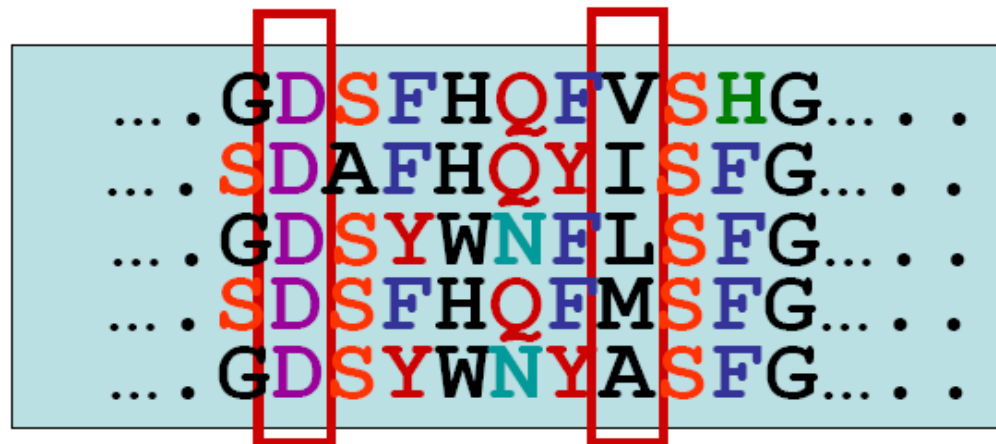
A sequence logo for the motif CAGGTAAGT. The logo shows the relative frequency of each nucleotide at each position. The positions are numbered from -3 to +6. The nucleotides are color-coded: C (red), A (green), G (blue), and T (yellow). The logo is connected by lines to the corresponding columns in the table below.

Pos	-3	-2	-1	+1	+2	+3	+4	+5	+6
A	0.3	0.6	0.1	0.0	0.0	0.4	0.7	0.1	0.1
C	0.4	0.1	0.0	0.0	0.0	0.1	0.1	0.1	0.2
G	0.2	0.2	0.8	1.0	0.0	0.4	0.1	0.8	0.2
T	0.1	0.1	0.1	0.0	1.0	0.1	0.1	0.0	0.5

PSSM: 思考与应用



- 可以根据BLOCK推导得到的PSSM进行数据库的搜索，发现包含该模式的新的蛋白质，并预测功能
- 需要思考的问题：
 - ✿ 根据PSSM如何计算新的序列？
 - ✿ PSSM中究竟包含着何等信息？



问题一：PSSM->发现



- ❑ 计算log-odds ratio/Odds ratio
- ❑ Do not miss: 性能检验!!!
- ❑ 结果需要计算 Sn , Sp , Ac & Mcc
- ❑ 需要计算Self-consistency, Leave-one-out validation & n -fold cross-validation

问题一：PSSM->发现



- ❑ 计算log-odds ratio/Odds ratio
- ❑ Do not miss: 性能检验!!!
- ❑ 结果需要计算 Sn , Sp , Ac & Mcc
- ❑ 需要计算Self-consistency, Leave-one-out validation & n -fold cross-validation

计算log-odds ratio



Pos	-3	-2	-1	+1	+2	+3	+4	+5	+6
A	0.3	0.6	0.1	0.0	0.0	0.4	0.7	0.1	0.1
C	0.4	0.1	0.0	0.0	0.0	0.1	0.1	0.1	0.2
G	0.2	0.2	0.8	1.0	0.0	0.4	0.1	0.8	0.2
T	0.1	0.1	0.1	0.0	1.0	0.1	0.1	0.0	0.5

$$P(S|+) = P_{-3}(S_1)P_{-2}(S_2)P_{-1}(S_3) \cdots P_5(S_8)P_6(S_9)$$

□ $P(S|+)$, 根据阳性训练数据计算出来的概率

Then, $P(S|-)$?



- ❑ 负样本/阴性数据的概率计算

- ❑ 计算方法:

 - ✿ A. DNA序列, 四种碱基出现的频率

 - ✿ B. 蛋白质序列, 20种氨基酸出现的频率

Odds Ratio



5' splice signal

Con:	C	A	G	...	G	T
Pos	-3	-2	-1	...	+5	+6
A	0.3	0.6	0.1	...	0.1	0.1
C	0.4	0.1	0.0	...	0.1	0.2
G	0.2	0.2	0.8	...	0.8	0.2
T	0.1	0.1	0.1	...	0.0	0.5

Background

Pos	Generic
A	0.25
C	0.25
G	0.25
T	0.25

$$S = S_1 S_2 S_3 S_4 S_5 S_6 S_7 S_8 S_9$$

$$\text{Odds Ratio: } R = \frac{P(S|+) = P_{-3}(S_1)P_{-2}(S_2)P_{-1}(S_3) \cdots P_5(S_8)P_6(S_9)}{P(S|-) = P_{bg}(S_1)P_{bg}(S_2)P_{bg}(S_3) \cdots P_{bg}(S_8)P_{bg}(S_9)}$$

Log-odds Ratio



$$\text{Odds Ratio: } R = \frac{P(S|+)}{P(S|-)} = \frac{P_{-3}(S_1)P_{-2}(S_2)P_{-1}(S_3) \cdots P_5(S_8)P_6(S_9)}{P_{bg}(S_1)P_{bg}(S_2)P_{bg}(S_3) \cdots P_{bg}(S_8)P_{bg}(S_9)}$$

$$= \prod_{k=1}^{k=9} P_{-4+k}(S_k) / P_{bg}(S_k)$$

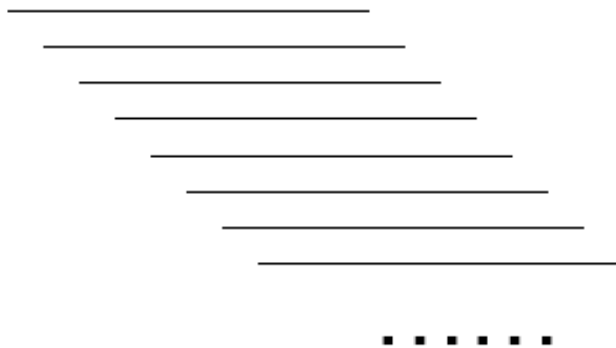
$$\text{Score } s = \log_2 R = \sum_{k=1}^{k=9} \log_2 (P_{-4+k}(S_k) / P_{bg}(S_k))$$

计算流程：滑动窗口



Slide WMM along sequence:

ttgacctagatgagatgtcgttcacttttactgagctacagaaaa



□ 设定域值；窗口宽度12 bp；依次打分预测

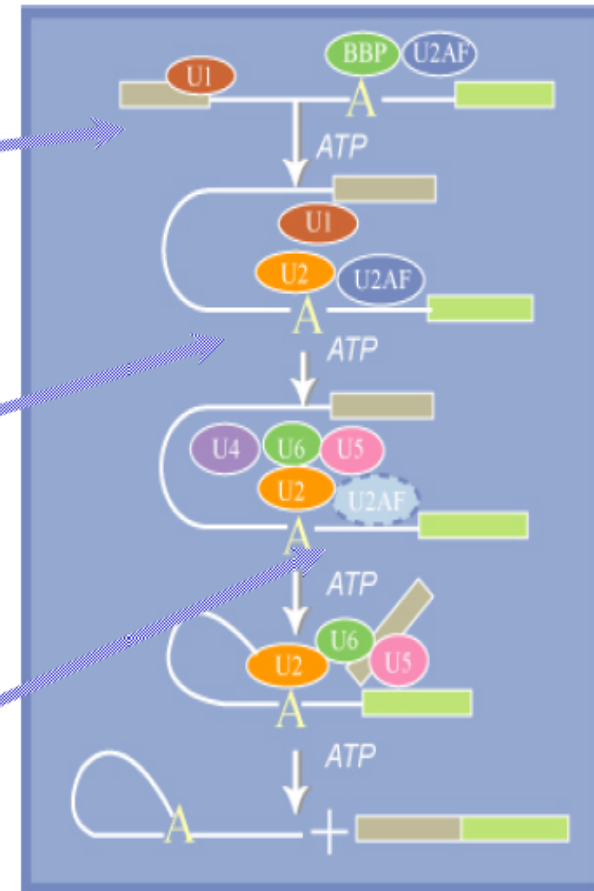
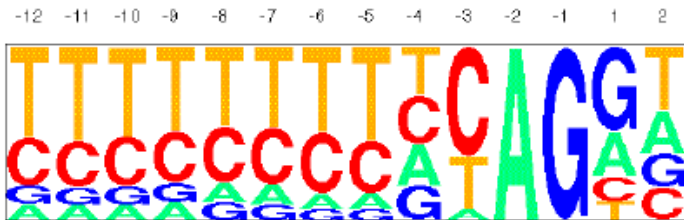
例：剪接模型 (Splicing)



branch site



3' splice site



计算log-odds ratio



5' splice signal

Con: C A G ... G T

Pos	-3	-2	-1	...	+5	+6
A	0.3	0.6	0.1	...	0.1	0.1
C	0.4	0.1	0.0	...	0.1	0.2
G	0.2	0.2	0.8	...	0.8	0.2
T	0.1	0.1	0.1	...	0.0	0.5

Background

Pos	Generic
A	0.25
C	0.25
G	0.25
T	0.25

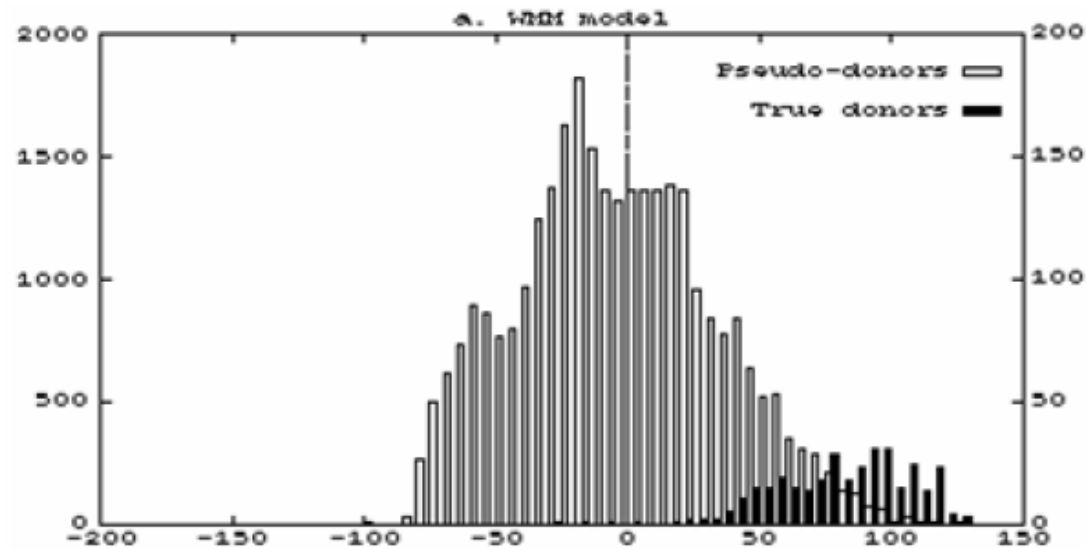
$$S = S_1 S_2 S_3 S_4 S_5 S_6 S_7 S_8 S_9$$

$$\text{Odds Ratio: } R = \frac{P(S|+) = P_{-3}(S_1)P_{-2}(S_2)P_{-1}(S_3) \cdots P_5(S_8)P_6(S_9)}{P(S|-) = P_{bg}(S_1)P_{bg}(S_2)P_{bg}(S_3) \cdots P_{bg}(S_8)P_{bg}(S_9)}$$

真实的打分情况：5' SS



"Decoy"
5'
Splice
Sites



True
5'
Splice
Sites

Sn:	<u>20%</u>	<u>50%</u>	<u>90%</u>
Sp:	50%	32%	7%

结果解释



- ❑ 细胞中的剪接机器（**Splicing machinery**）可能识别其他的，不包括在训练数据中的模式
- ❑ **PSSM模型不能很好的反映真实的5' SS的识别情况**
- ❑ 两种可能：**either or both**

Log-odds ratio vs. 贝叶斯



- X-box序列, 1000个4 bp的DNA序列有100个为真实的X-box。分布如下:

	第一位	第二位	第三位	第四位
A	70%	10%	1%	5%
T	10%	10%	97%	5%
C	10%	70%	1%	5%
G	10%	10%	1%	85%

$$P(X\text{-box})=0.1$$

$$P(X - box | X_1 X_2 X_3 X_4) = \frac{P(X - box) \prod_i q_{xi}^{x-box}}{P(X - box) \prod_i q_{xi}^{x-box} + P(nonX - box) \prod_i q_{xi}^{nonX-box}}$$

Log-odds ratio vs. 贝叶斯 (2)



- ❑ 贝叶斯方法：必须估计真实的与错误的数量上的比例；另外，假设在未知数据中，(+)与(-)之间的比例不变
- ❑ Log-odds ratio: 不需要知道(+)与(-)之间的比例；观测到的数据，其为(+)与(-)的比值
- ❑ 域值的确定：都需要计算 S_n , S_p , Ac & MCC

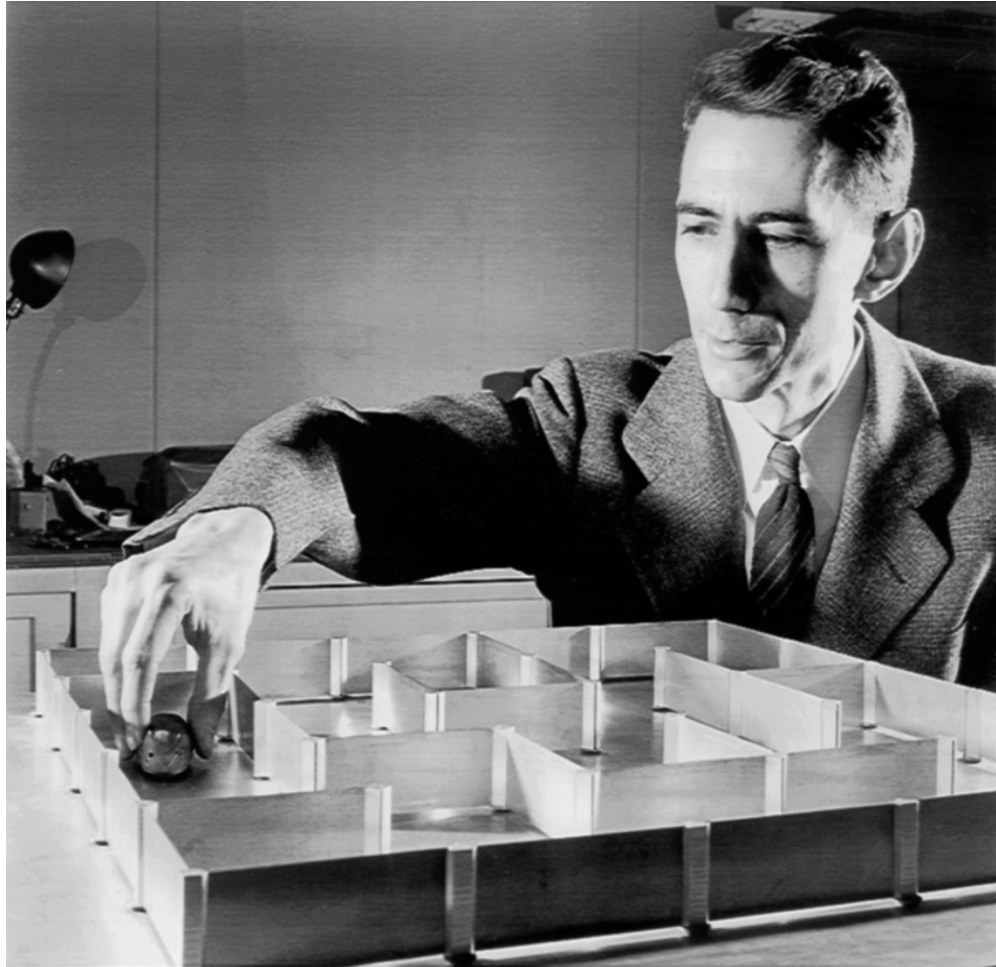
问题二：PSSM->信息？



- ❑ PSSM/motif/domain/BLOCK: 每一个位置上究竟包含了什么样的信息？
- ❑ 对于同一个motif/PSSM: 有些位点较其他位点提供更多的信息，why？
- ❑ 如何定量化“信息”？



信息论：Claude Shannon

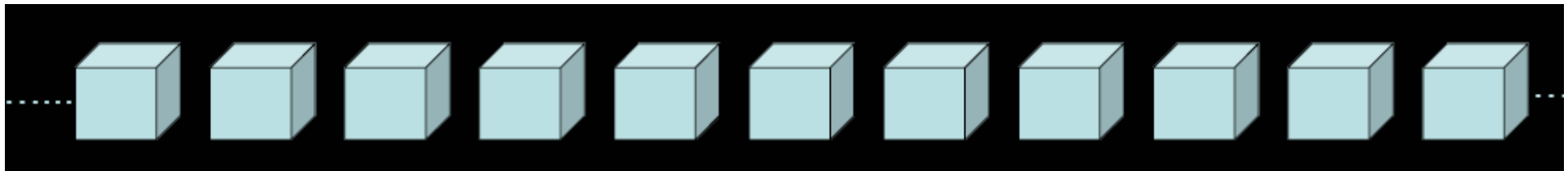


□ 信息论的奠基人

1,048,576个盒子：Yes/No?



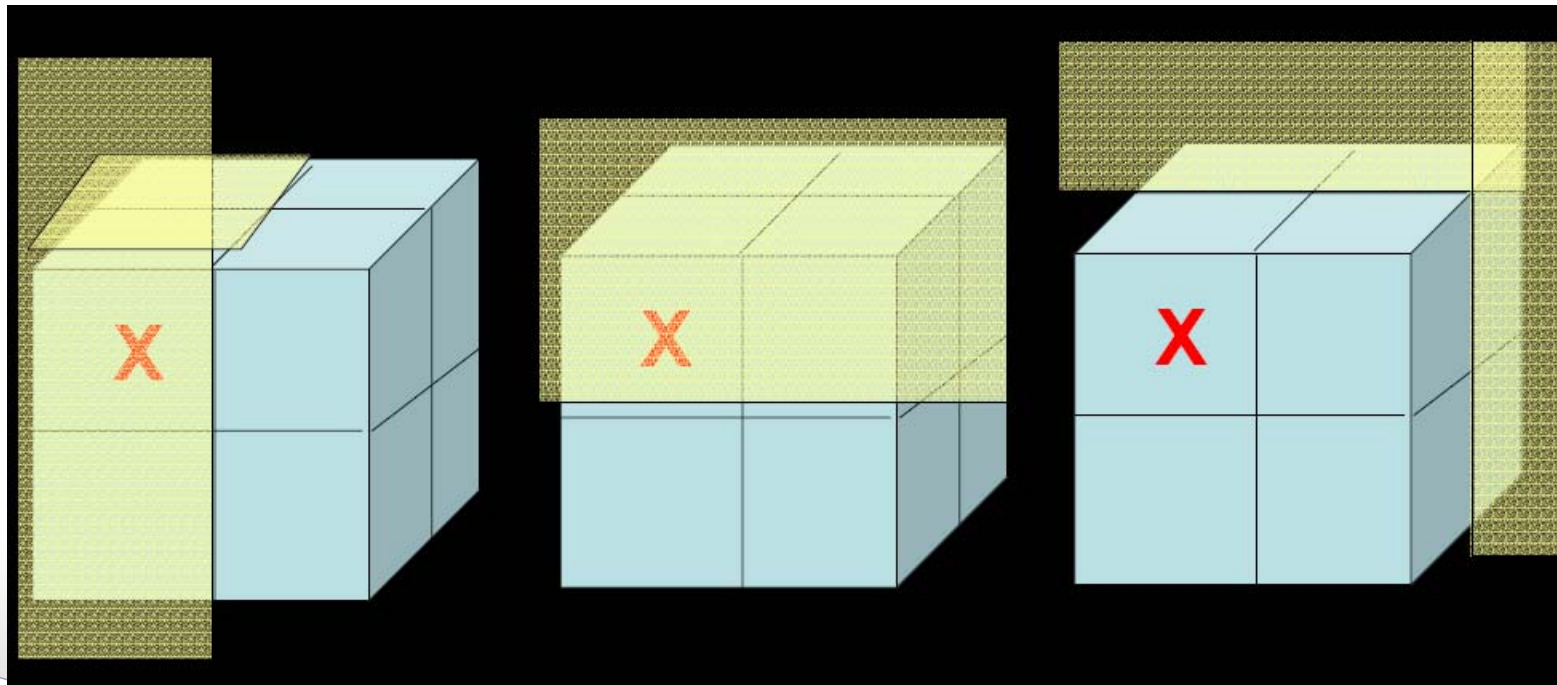
- ❑ 随机将10000 RMB的支票放入1,048,576个盒子之一
- ❑ Play 20 questions: yes/no



8个盒子



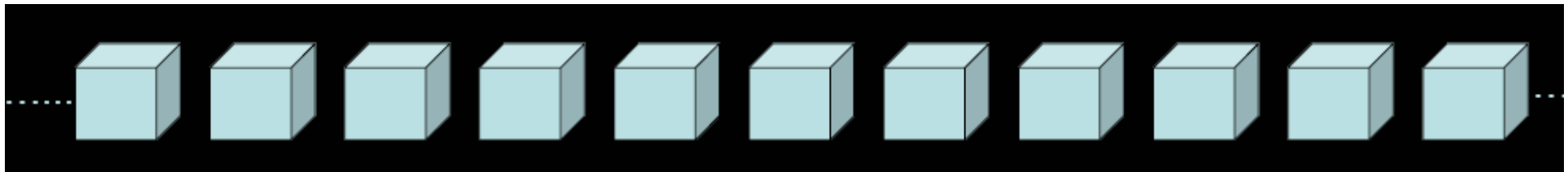
- ❑ 最少问多少个yes/no的问题能够定位支票?
- ❑ Answer: $\log_2 8 = 3$



1,048,576个盒子：Yes/No?



- ❑ 随机将10000 RMB的支票放入1,048,576个盒子之一
- ❑ Play 20 questions: yes/no



❑ $2^{20} = 1,048,576$

信息论



- $2^b = M$; b 为bit (binary digit) 信息
- M : 所有概率事件的总数; 因此:
- $b = \log_2(M)$; $\Rightarrow b = -\log_2(1/M) \Rightarrow b = -\log_2(P)$; 所有事件概率相同, 则 $P=1/M$
- 例: 对于某一个motif的一个位置上, 可能存在20种氨基酸, 且概率相等, 则 $P=1/20$
- 香农熵: $b = -\log_2(1/20) = 4.32 \text{ bits}$

信息论 (2)



□ 若概率不等同，如何处理？

□ 定义 $u_i = -\log_2(P_i)$

...	G	D	S	F	H	Q	E	V	S	H	G	...
...	S	D	A	F	H	Q	Y	I	S	F	G	...
...	G	D	S	Y	W	N	E	L	S	F	G	...
...	S	D	S	F	H	Q	E	M	S	F	G	...
...	L	D	S	Y	W	N	Y	A	S	F	G	...

$$\text{平均熵值} = \frac{u_V + u_I + u_L + u_M + u_A}{N}$$

N: 全部序列的数目

香农熵的平均值 =

$$\frac{\sum_{i=1}^{20} N_i u_i}{N}$$

N_i : 在该位置上为氨基酸i的序列的数目

信息论 (3)



$$\frac{\sum_{i=1}^{20} N_i u_i}{N} \Rightarrow \sum_{i=1}^{20} \frac{N_i u_i}{N}$$

□ 上式中， $N_i/N=P_i$ ；因此，上式可转化为： $\sum_{i=1}^{20} P_i u_i$

□ 因此，香农熵公式为：

$$H = - \sum_{i=1}^{20} P_i (\log_2 P_i)$$

信息论：意义？



- 香农的信息熵公式：
$$H = - \sum_{i=1}^{20} P_i (\log_2 P_i)$$
- H为每个位置上的“香农熵”
- 香农熵：不确定性！
- 在每一个位置上，各种氨基酸出现的不确定性

信息论 (4)



... G D S F H Q E V S H G ...
... S D A F H Q Y I S F G ...
... G D S Y W N E L S F G ...
... S D S F H Q F M S F G ...
... L D S Y W N Y A S F G ...

□ $P(D)=1$, 因此, $H = -1 \cdot \log_2(1) = -1 \cdot 0 = 0$

无不确定性

□ $P(V) = P(I) = P(L) = P(M) = P(A) = 1/5$;

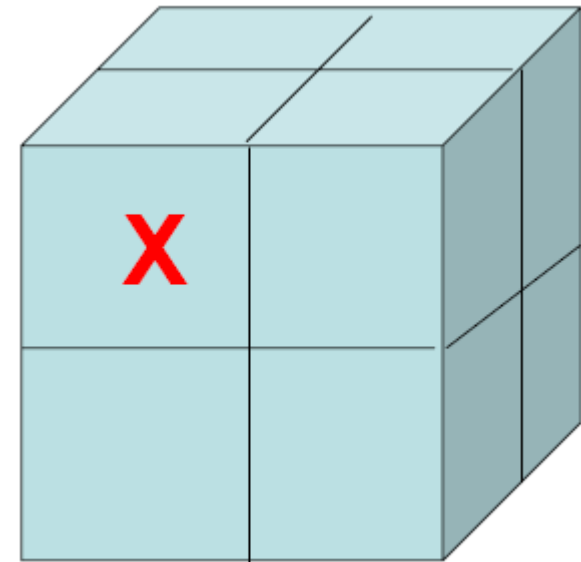
□ $H = -(1/5) \cdot \log_2(1/5) - (1/5) \cdot \log_2(1/5) -$
 $(1/5) \cdot \log_2(1/5) - (1/5) \cdot \log_2(1/5) -$
 $(1/5) \cdot \log_2(1/5) = 2.32 \text{ bit}$

高不确定性

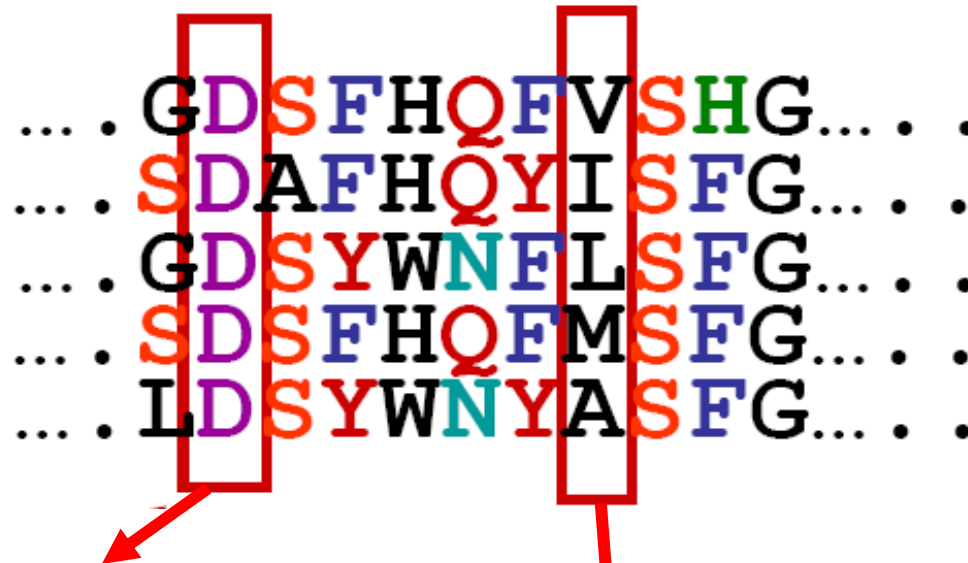
不确定性 -> 信息



- 盒子模型；
- 假设只能回答两个问题；则
 - ✿ A. 回答问题之前，不确定性为3 bits
 - ✿ B. 回答问题之后，不确定性为1 bit
- 获得信息R：
- $R = H_{\text{before}} - H_{\text{after}} = 3 - 1 = 2$ bits



不确定性 -> 信息 (2)



- 假设，所有氨基酸出现的频率是相等的；则
- $H_{\text{before}} = 4.32$;
- $H_{\text{after}} = 0$;
- Motif在该位置的信息量为：4.32 bits

- $H_{\text{before}} = 4.32$;
- $H_{\text{after}} = 2.32$;
- Motif在该位置的信息量为：2 bits

模体发现: Gibbs Sampler



- Gibbs Sampler是一种Monte-Carlo类的方法，对于输入序列，找到一个最大的似然函数
- 对于序列s，且在位置A有一个motif的似然函数，定义如下：

weight matrix background
freq. vector

$$P(s, A | \Theta, \theta_B) =$$
$$\theta_{B,a} \times \dots \times \theta_{B,g} \times \Theta_{1,t} \times \Theta_{2,a} \times \dots \times \Theta_{8,c} \times \theta_{B,t} \times \dots \times \theta_{B,t}$$

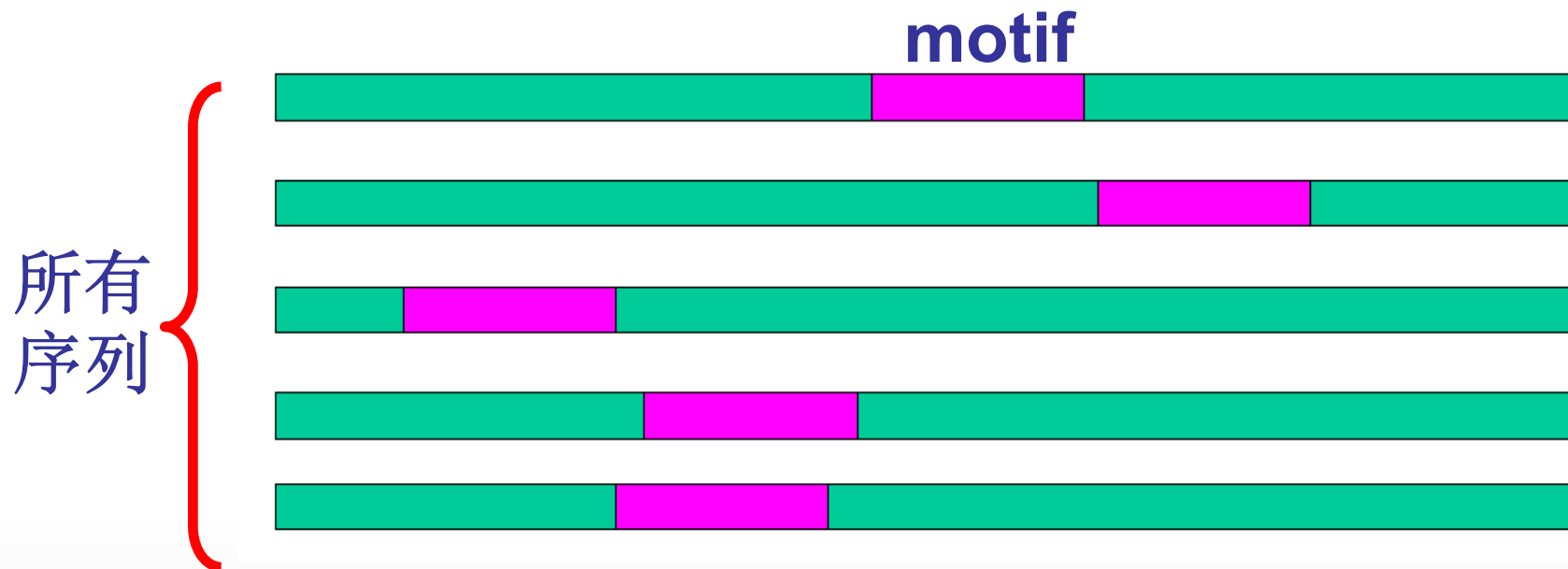
$s = \text{"actactgtatcgtactgactgattaggccatgactgcat"}$

Motif location A

Gibbs Sampling 算法 (1)



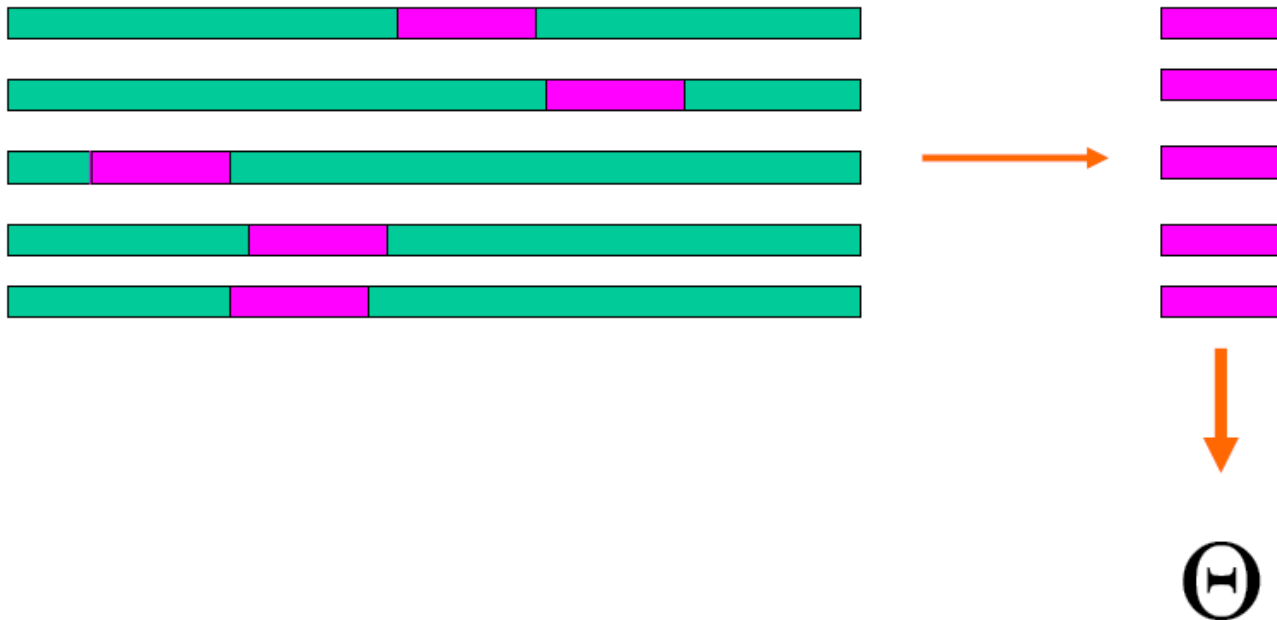
- 从每条序列上随机的抽取一段序列，序列长度固定



Gibbs Sampling 算法 (2)



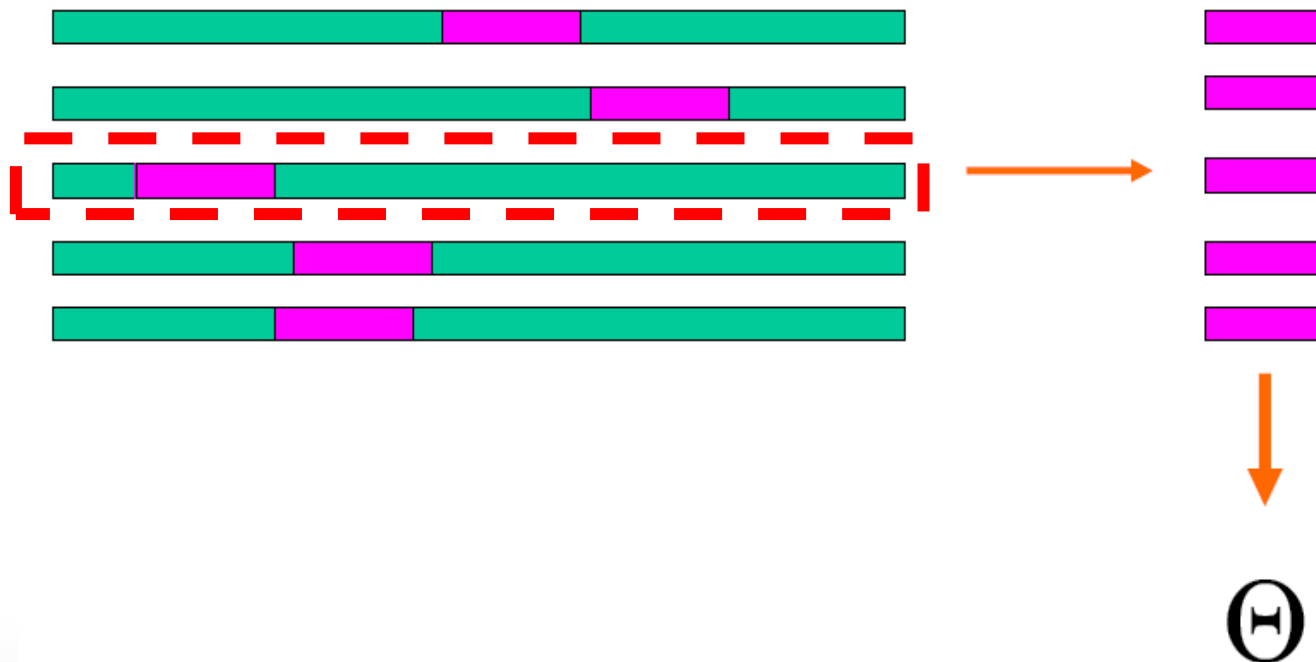
□ 构建PSSM/权重矩阵



Gibbs Sampling 算法 (3)



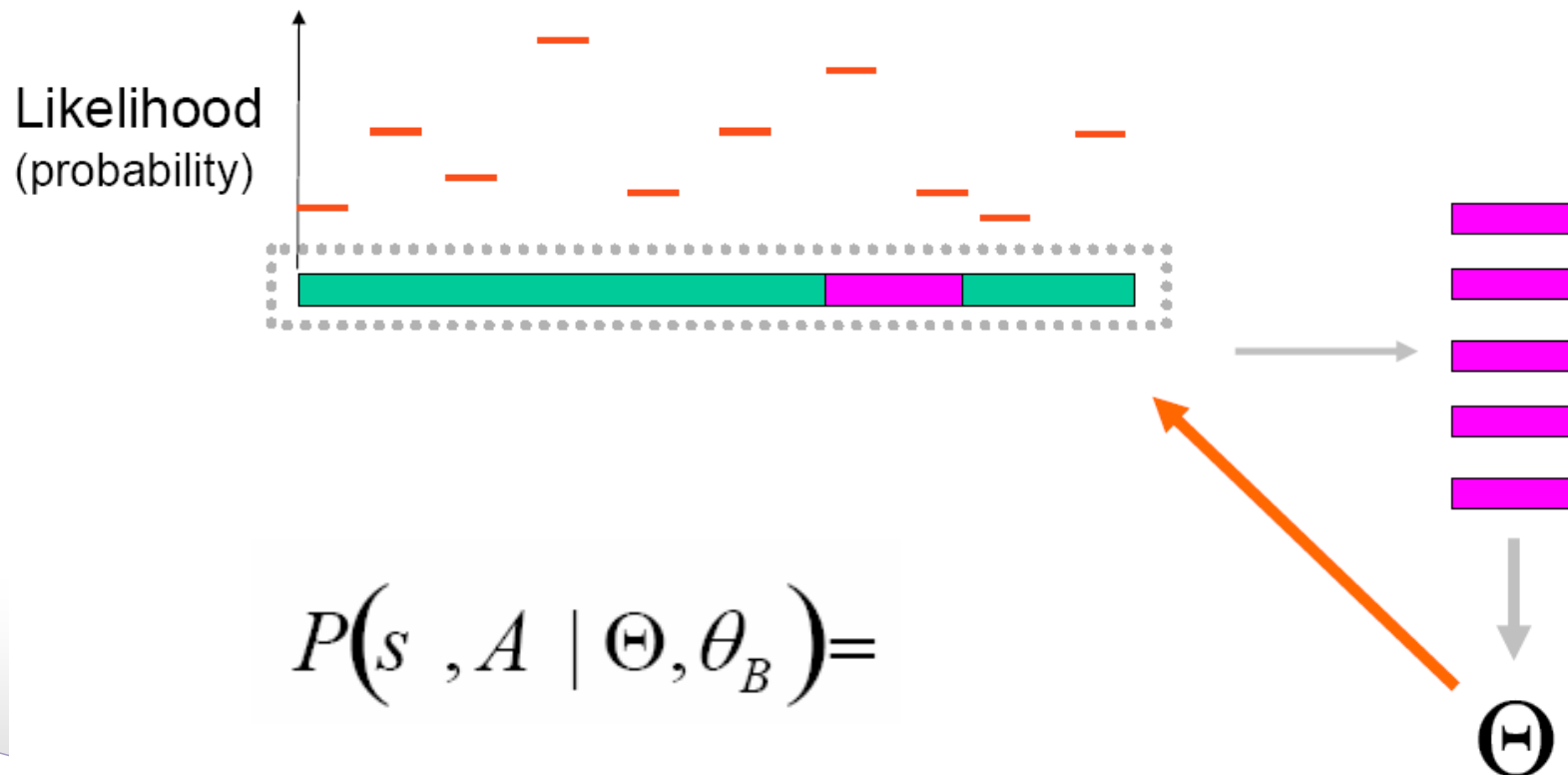
□ 随机挑选一条序列



Gibbs Sampling 算法 (4)



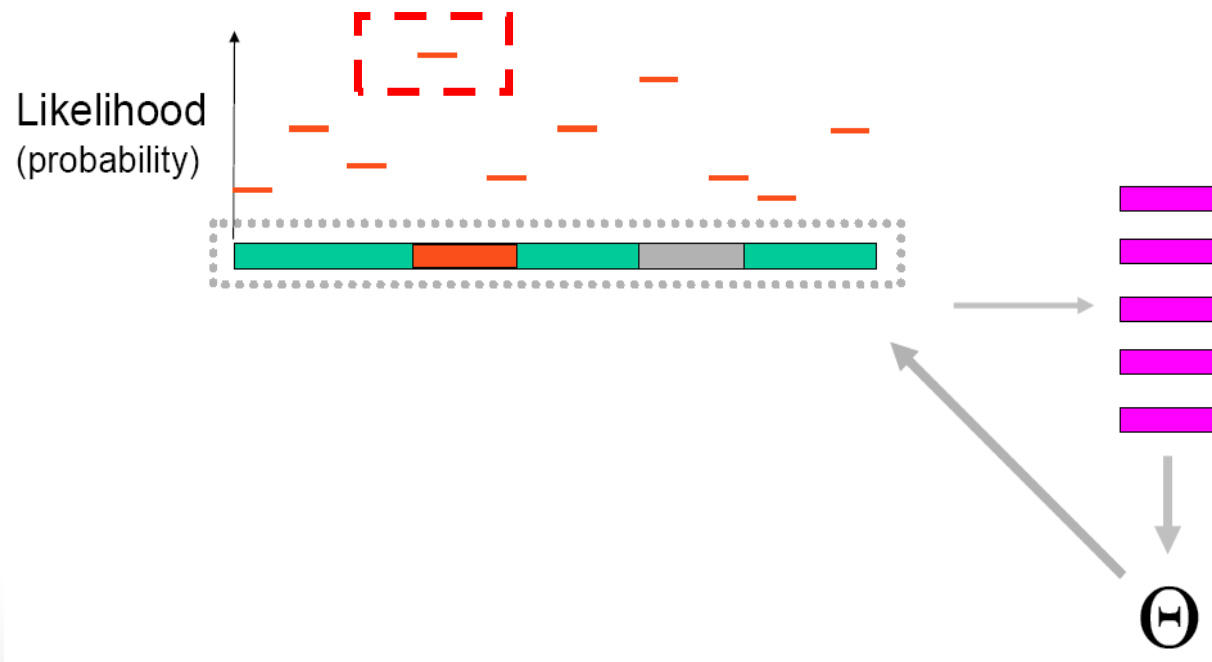
- 用构建好的PSSM对该序列上所有可能的 motif 进行打分 (窗口滑动, 每次1个氨基酸或者碱基)



Gibbs Sampling 算法 (5)



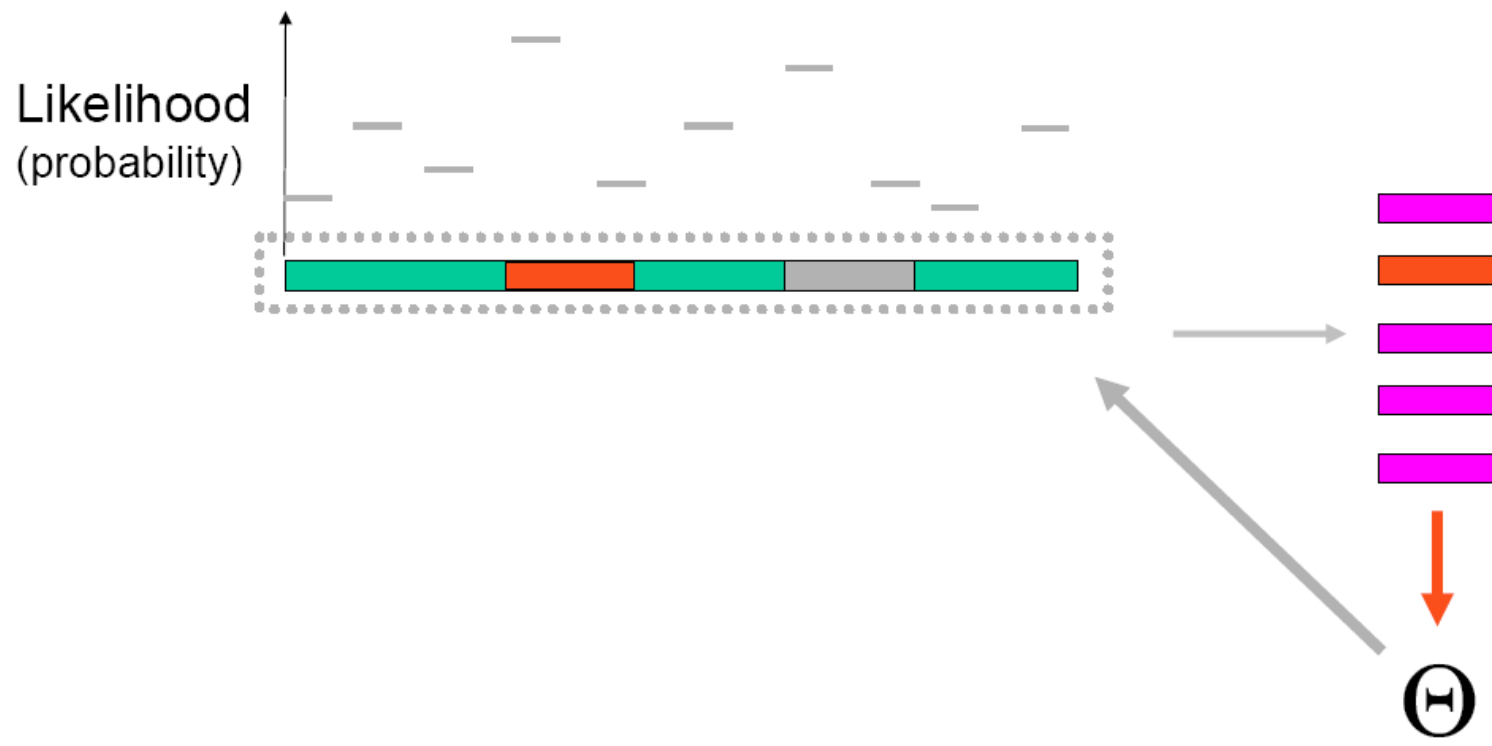
- 根据似然性的计算，得到似然值最大的模体，即新的motif



Gibbs Sampling 算法 (6)



□ 更新PSSM矩阵



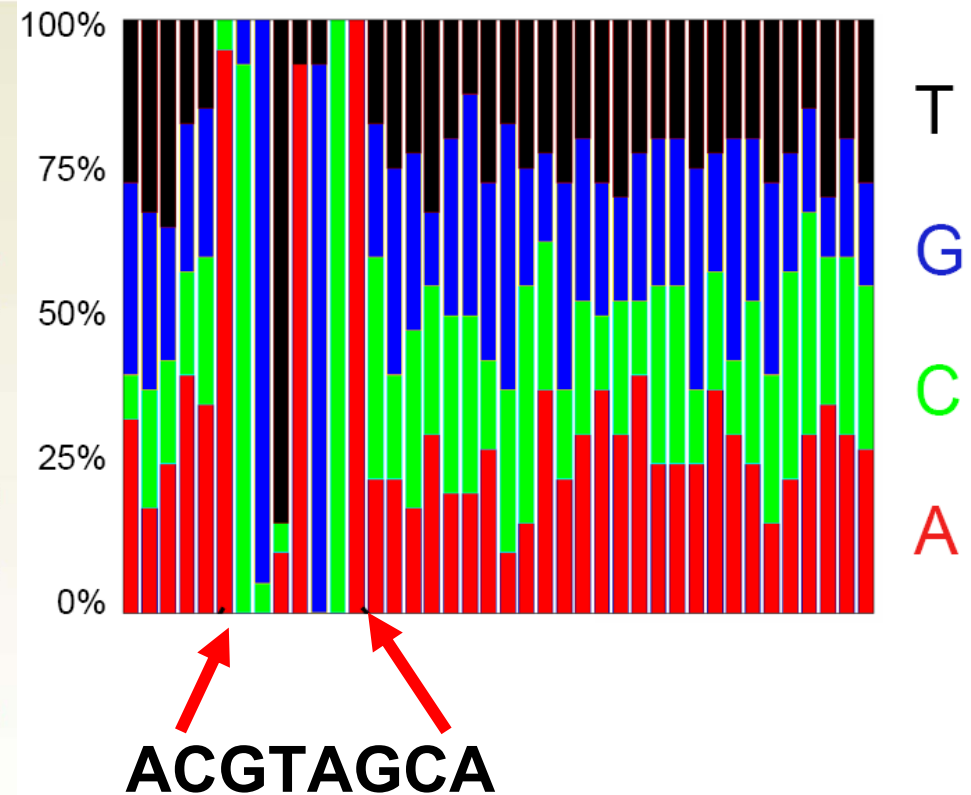
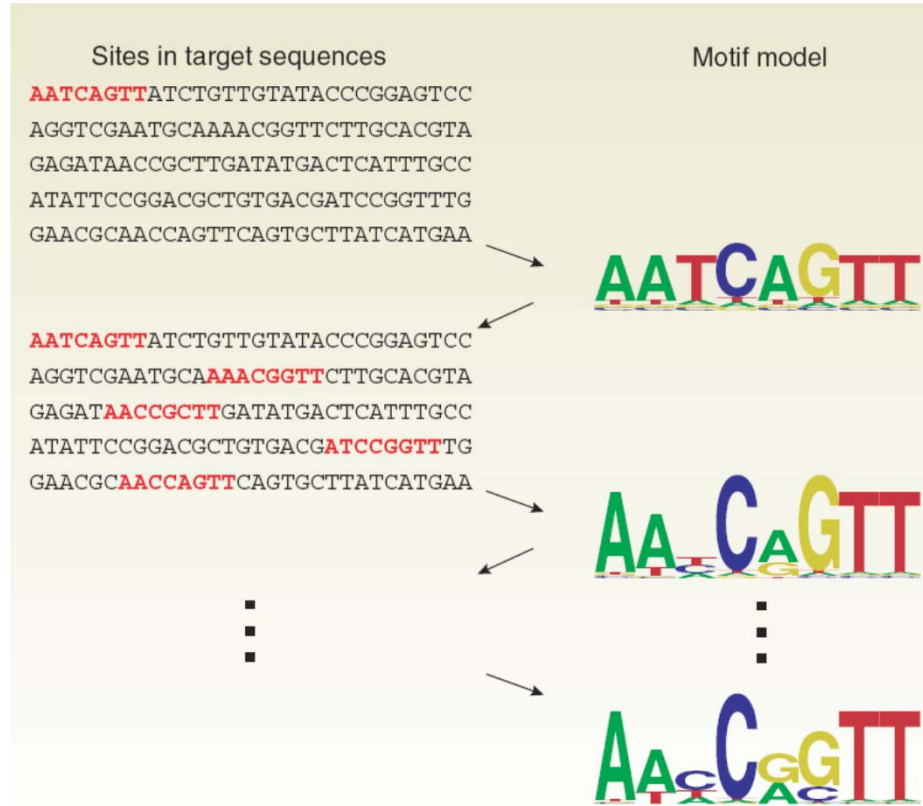
Gibbs Sampling 算法 (7)



- 反复迭代计算，直到似然性结果与PSSM不再发生变化



Strong Motif



Gibbs Sampler: 总结



- ❑ 模体发现的一种随机算法 (Monte Carlo)
- ❑ 寻找次优解的算法
- ❑ 根据PSSM/WMM对随机抽取的序列进行打分来调整采样，直到结果收敛
- ❑ 不能够保证每次运算的结果一致：需要多运算几次，并进行比较
- ❑ 对蛋白质、DNA、RNA序列模体的发现有帮助

期望最大化算法

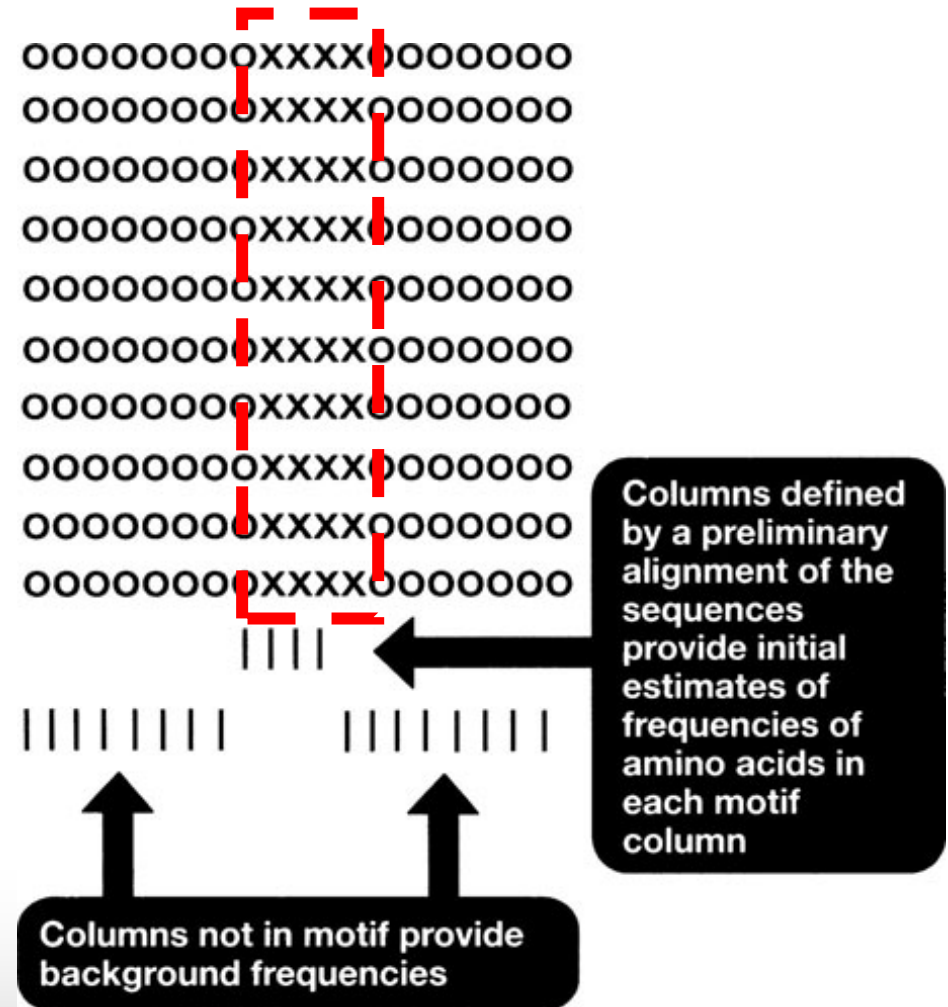


- Expectation Maximization Algorithm
- 已开发工具: Multiple EM for Motif Elicitation (MEME)
 - ✿ <http://meme-suite.org/>
- motif大致的位置与长度是确定的
- 重点: 确定motif在每条序列上的起始位置
- 分为两步:
 - ✿ E step: 估计motif起始位置的期望最大化
 - ✿ M step: motif似然性的期望最大化

期望最大化算法 (2)



- ❑ 假设10条序列，长度20个碱基
- ❑ 进行多序列，大致确定motif的位置
- ❑ 待找motif长度为5个碱基



Motif的概率 vs. 背景概率



- ❑ 计算motif中每个位置的碱基的概率分布
- ❑ 背景概率：根据剩下的序列计算四种碱基的概率分布

	Background	Site column 1	Site column 2	...
G	0.27	0.4	0.1	...
C	0.25	0.4	0.1	...
A	0.25	0.2	0.1	...
T	0.23	0.2	0.7	...
	1.00	1.0	1.0	



似然性概率值的计算

$$\Pr(X_i | Z_{ij} = 1, p) = \underbrace{\prod_{k=1}^{j-1} P_{ck0}}_{\text{before motif}} \underbrace{\prod_{k=j}^{j+w-1} P_{ck, k-j+1}}_{\text{motif}} \underbrace{\prod_{k=j+w}^L P_{ck0}}_{\text{after motif}}$$

其中, X_i 指的是第 i 个序列; 如果模体起始于序列 i 的位置 j , 那么就有 $Z_{ij} = 1$; C_k 指的是位于序列 i 中位置 k 的字母。

背景概率

例如: 假设 $X_i = G C A G T A G$

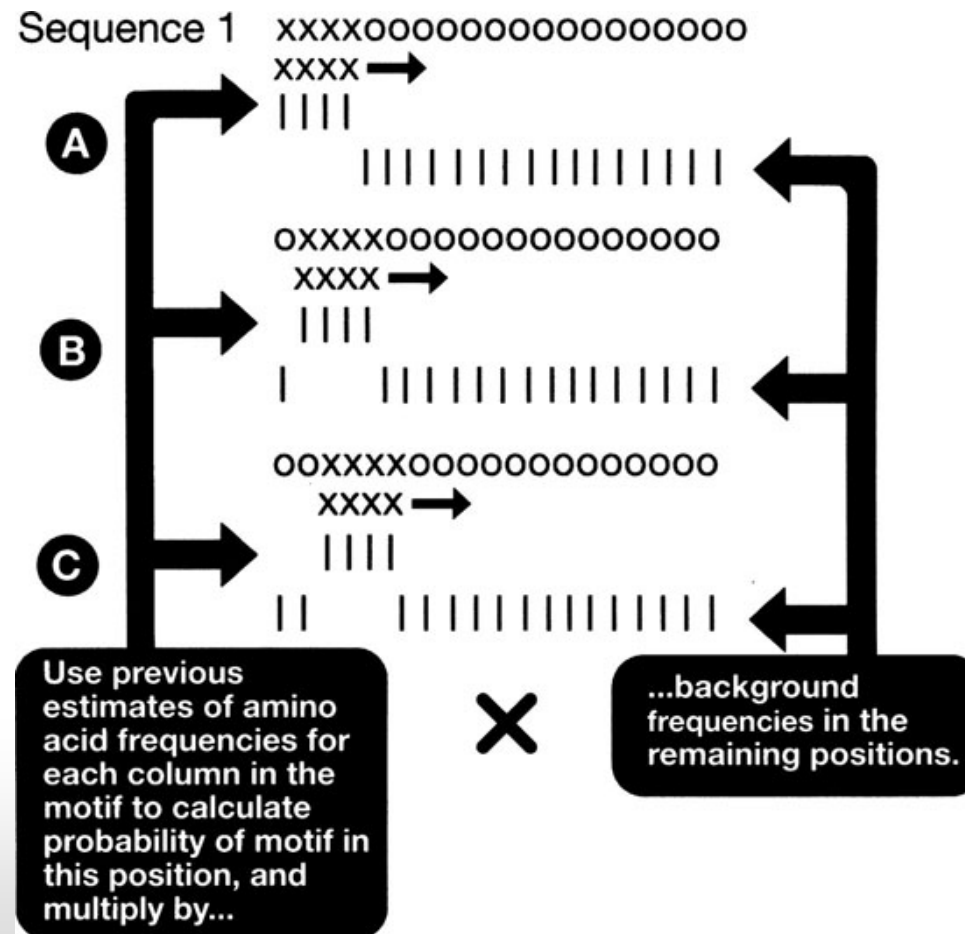
	0	1	2	3
A	0.25	0.1	0.5	0.2
$P = C$	0.25	0.4	0.2	0.1
G	0.25	0.3	0.1	0.6
T	0.25	0.2	0.2	0.1

那么 $\Pr = (X_i | Z_{i3} = 1, p) = P_{G,1} \times P_{C,0} \times P_{T,1},$
 $P_{G,2} \times P_{T,3} \times P_{A,0} \times P_{G,0} = 0.25 \times 0.25 \times 0.2 \times 0.1 \times$
 $0.1 \times 0.25 \times 0.25$

似然性概率值的计算 (2)



- 计算每条序列，在不同的起始位置，其似然性的概率值





E step: 起始位置估计

□ Z值: motif在不同位置起始的几率值

$$Z_{ij}^{(t)} = \frac{\Pr(X_i | Z_{ij} = 1, P^{(t)}) \Pr(Z_{ij} = 1)}{\sum_{k=1}^{l-w+1} \Pr(X_i | Z_{ik} = 1, P^{(t)}) \Pr(Z_{ik} = 1)}$$

□ 假设, motif在任意位置起始的概率相同, 则

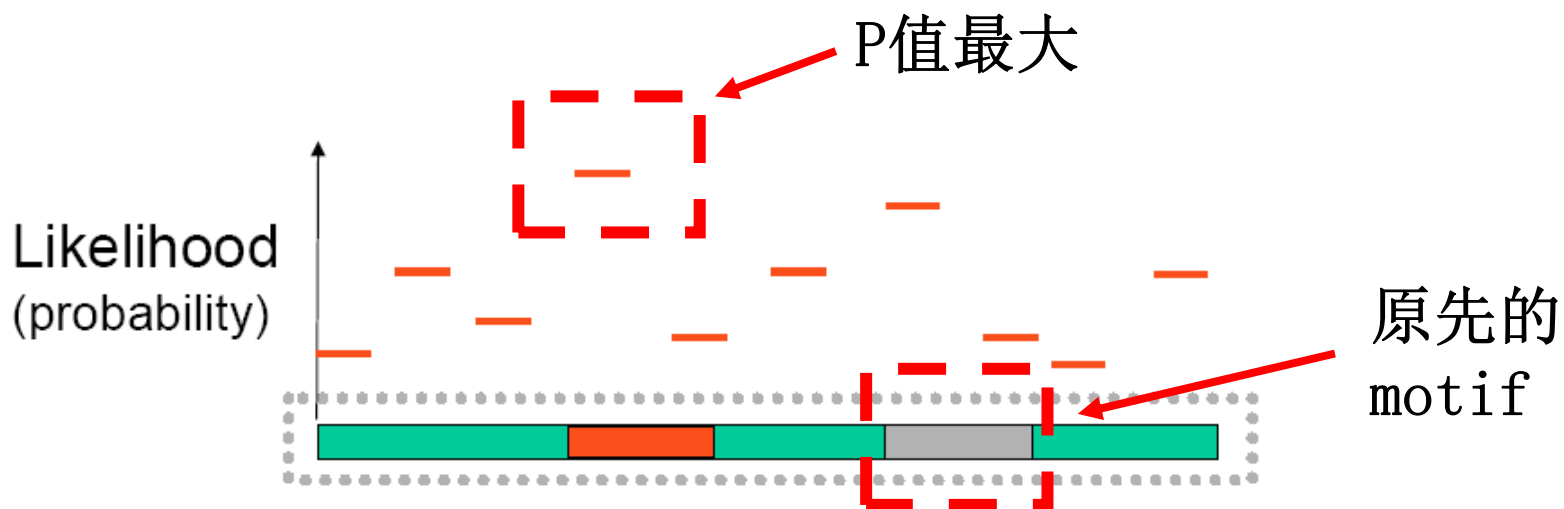
$$Z_{ij}^{(t)} = \frac{\Pr(X_i | Z_{ij} = 1, P^{(t)})}{\sum_{k=1}^{l-w+1} \Pr(X_i | Z_{ik} = 1, P^{(t)})}$$

□ Z值最大化, 即为“最可能的起始位置”

M step: P值最大化



- 根据选择的最大Z值，重新计算矩阵，并计算P值最大的motif



EM算法：迭代



- a. 给定长度参数 W 和一组序列;
- b. 设定 P 的初始值;
- c. Do
 - 根据 P 反复进行 Z 的估计; (E - 步骤)
 - 根据 Z 反复进行 P 的估计。 (M - 步骤)
- d. 直到 P 的变化小于 ϵ 为止;
- e. 返回 P, Z 的值。

Gibbs & EM: 总结



- ❑ 基本假设：所有序列都拥有，且仅拥有一个motif
- ❑ 估算两个关联的函数：Gibbs (WMM & motif的似然性)，EM (motif起始位置，Z 值 & motif的似然性)
- ❑ 利用两个函数的其中之一修正另一个，采取迭代/反复计算的方法，使结果收敛
- ❑ 不保证得到的结果为最优，近似算法