

Computational Biology

Lecture 10



Saad Mehmeh

Hidden Markov Model

A Hidden Markov Model HMM is defines as:

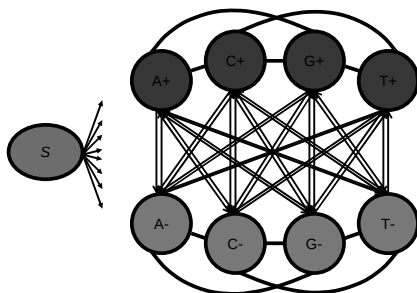
- decouple states from symbols
- A set of *hidden* states
 - [transitional probabilities]
For each pair of states i and j ,
a transition probability a_{ij} .
 - $\sum_j a_{ij} = 1$
 - An alphabet of symbols Σ
 - [emission probabilities]
For each state k , and symbol b
 $e_k(b) = p(x_i = b \mid \pi_i = k)$
[now we use variable π for states
and variable x for symbols]
 - $\sum_{b \in \Sigma} e_k(b) = 1$ for each state k

Markov property: $p(\pi_n = j \mid x_0 \dots x_m, \pi_0 \dots \pi_{m-1}, \pi_m = i) =$
 $p(\pi_n = j \mid \pi_m = i) \quad m < n$
if $m = n - 1$, then this is a_{ij}



Saad Mehmeh

HMM for CpG islands



Saad Mehmeh

Questions with HMMs

- **Evaluation:** given x , what is the probability $p(x)$ that it was produced by the model?
- **Decoding:** given x , what is the most probable path that produces x in the model?
- **Learning:** given x , what are the parameters (transitional probabilities and emission probabilities) of the model that maximize $p(x)$.



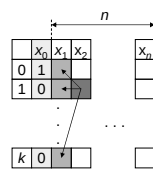
Saad Mehmeh

Viterbi decoding algorithm

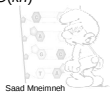
- Initialization
 $v_0(0) = 1$, $v_k(0) = 0$ for $k > 0$

- Main iteration
 for $i = 1 \dots n$
 $v_i(i) = e_i(x_i) \cdot \max_k (v_{k0}(i-1) \cdot a_{ki})$
 $ptr_i(i) = \operatorname{argmax}_k (v_{k0}(i-1) \cdot a_{ki})$

- Termination
 $p(x, \pi^*) = \max_k (v_k(n) a_{k0})$



Time = $O(k^2n)$
 Space = $O(kn)$



Saad Mehmeh

Computing $p(x)$

- Before, $p(x) = a_{sx1} \prod_{i=2 \dots n} a_{x^{i-1} x^i}$
- Now, $p(x) = \sum_{\pi} p(x, \pi)$
- Enumerating all π is exponential!
- Use Viterbi, same as before, but change **max** to **Σ**



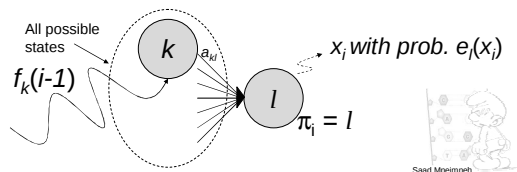
Saad Mehmeh

Forward *evaluation* algorithm

- Let $f_i(l) = p(x_1 \dots x_i, \pi_i = l)$

- Then,

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



Saad Mehmeh

Derivation

$$\begin{aligned} f_i(l) &= \sum_{x_1 \dots x_{i-1}} p(x_1 \dots x_i, \pi_1 \dots \pi_{i-1}, \pi_i = l) \\ &= \sum_{x_1 \dots x_{i-1}} p(x_i, \pi_i = l, x_1 \dots x_{i-1}, \pi_1 \dots \pi_{i-1}) \\ &= \sum_{x_1 \dots x_{i-1}} p(x_i, \pi_i = l \mid x_1 \dots x_{i-1}, \pi_1 \dots \pi_{i-1}) \cdot p(x_1 \dots x_{i-1}, \pi_1 \dots \pi_{i-1}) \\ &= \sum_{x_1 \dots x_{i-1}} p(x_i, \pi_i = l \mid \pi_{i-1}) \cdot p(x_1 \dots x_{i-1}, \pi_1 \dots \pi_{i-1}) \\ &= \sum_{x_1 \dots x_{i-2}, k} p(x_i, \pi_i = l \mid \pi_{i-1} = k) \cdot p(x_1 \dots x_{i-1}, \pi_1 \dots \pi_{i-2}, \pi_{i-1} = k) \\ &= \sum_{x_1 \dots x_{i-2}, k} e_l(x_i) a_{kl} \cdot p(x_1 \dots x_{i-1}, \pi_1 \dots \pi_{i-2}, \pi_{i-1} = k) \\ &= \sum_k \sum_{x_1 \dots x_{i-2}} e_l(x_i) a_{kl} \cdot p(x_1 \dots x_{i-1}, \pi_1 \dots \pi_{i-2}, \pi_{i-1} = k) \\ &= \sum_k \sum_{x_1 \dots x_{i-2}} p(x_1 \dots x_{i-1}, \pi_1 \dots \pi_{i-2}, \pi_{i-1} = k) e_l(x_i) a_{kl} \\ &= \sum_k f_k(i-1) e_l(x_i) a_{kl} = e_l(x_i) \sum_k f_k(i-1) a_{kl} \end{aligned}$$

Saad Mehmeh

Forward *evaluation* algorithm

- Initialization

$$f_0(0) = 1, f_k(0) = 0 \text{ for } k > 0$$

- Main iteration

for $i = 1 \dots n$

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

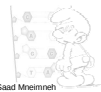
- Termination

$$p(x) = \sum_k f_k(n) a_{k0}$$

Saad Mehmeh

Problem of small numbers

- In Viterbi and Forward algorithm we multiply probabilities $\rightarrow$ numbers will soon be very small and we lose precision.
- Use log space $\rightarrow$ addition instead of multiplication.



Saad Mehinneh

Log space Viterbi

$$v_i(i) = e_i(x_i) \cdot \max_k (v_k(i-1) \cdot a_{ki})$$

$$\text{Let } V_i(i) = \log v_i(i)$$

$$\begin{aligned} V_i(i) &= \log [e_i(x_i) \cdot \max_k (v_k(i-1) \cdot a_{ki})] \\ &= \log e_i(x_i) + \log [\max_k (v_k(i-1) \cdot a_{ki})] \\ &= \log e_i(x_i) + \max_k \log [(v_k(i-1) \cdot a_{ki})] \\ &= \log e_i(x_i) + \max_k (V_k(i-1) + \log a_{ki}) \end{aligned}$$



Saad Mehinneh

Log space Forward

$$f_i(i) = e_i(x_i) \cdot \sum_k (f_k(i-1) \cdot a_{ki})$$

$$\text{Let } F_i(i) = \log f_i(i)$$

$$\begin{aligned} F_i(i) &= \log [e_i(x_i) \cdot \sum_k (f_k(i-1) \cdot a_{ki})] \\ &= \log e_i(x_i) + \log \sum_k (f_k(i-1) \cdot a_{ki}) \\ &\neq \log e_i(x_i) + \sum_k \log [(f_k(i-1) \cdot a_{ki})] \end{aligned}$$

$$= \log e_i(x_i) + \log \sum_k e^{(F_k(i-1) + \log a_{ki})}$$



Saad Mehinneh

Back to the most probable path

- The Viterbi algorithm finds it!
- The most probable path might not be the most appropriate basis for judgment.
- We might want, for instance, the most probable state for an observation x_i .
- More generally, we are interested in $p(\pi_i = k \mid x)$



Saad Mehinneh

Computing $p(\pi_i = k \mid x)$

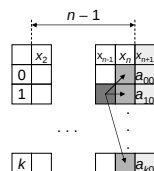
- $p(\pi_i = k \mid x) = p(x, \pi_i = k) / p(x)$
- I know how to compute $p(x)$: forward alg.
- $p(x, \pi_i = k)$
 $= p(x_1 \dots x_i, \pi_i = k) \cdot p(x_{i+1} \dots x_n \mid x_1 \dots x_i, \pi_i = k)$
 $= p(x_1 \dots x_i, \pi_i = k) \cdot p(x_{i+1} \dots x_n \mid \pi_i = k)$
 $= f_k(i) \cdot b_k(i)$



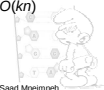
Saad Mehinneh

Backward *evaluation* algorithm

- Initialization
 $b_k(n) = a_{k0}$ for all k
- Main iteration
 for $i = n - 1 \dots 1$
 $b_k(i) = \sum_l a_{kl} e_l(x_{i+1}) b_l(i+1)$
- Termination
 $p(x) = \sum_l a_{l0} e_l(x_1) b_l(1)$



Time = $O(k^2n)$
 Space = $O(kn)$



Saad Mehinneh

Learning (training the HMM)

- Let θ be the parameters of the HMM (transition probabilities and emission probabilities, the a's and e's)
- Given independent sequences $x^1, \dots, x^n$, we would like to find θ that will maximize:

$$\log p(x^1 \dots x^n | \theta) = \sum_{j=1..n} \log p(x^j | \theta)$$

This is called the maximum likelihood parameters.



Saad Meimneh

State sequence is known

- Assume the path for each x^i is known
 - For instance, we have sequences in which CpG islands are already labeled
- Paths are known, let
 - A_{kl} = number of transitions from k to l
 - $E_k(b)$ = number of times b emitted in state k
- The maximum likelihood parameters are given by:
 - $a_{kl} = A_{kl} / \sum_l A_{kl}$
 - $e_k(b) = E_k(b) / \sum_b E_k(b)$



Saad Meimneh

Maximum likelihood from counts

- Assume we have a sequence of independent observations $x_1 \dots x_n$ and that we count n_i occurrences of outcome i , $i=1 \dots k$.
- Let θ_i = probability of i .
- Then $\theta^{\text{ML}} = \{\theta_i = n_i/n, i=1 \dots k\}$ is the maximum likelihood solution for θ .
- Consider any other θ . We want to show that
$$p(x | \theta^{\text{ML}}) > p(x | \theta)$$



Saad Meimneh

Proof

$$\begin{aligned}\log \frac{p(x|\theta^{ML})}{p(x|\theta)} &= \log \frac{\prod_i (\theta_i^{ML})^{n_i}}{\prod_i \theta_i^{n_i}} \\ &= \sum_i n_i \log \frac{\theta_i^{ML}}{\theta_i} \\ &= n \sum_i \theta_i^{ML} \log \frac{\theta_i^{ML}}{\theta_i} > 0\end{aligned}$$

The last summation is the relative entropy of θ^{ML} and θ which is always positive and 0 iff $\theta^{ML} = \theta$ (from information theory)



Saad Mehmeh

Some problems

- Maximum likelihood are vulnerable to overfitting if insufficient data.
- For instance, if a state k was never used in the set of training sequences, then
 - $a_{kl} = 0$ for all l
 - $e_k(b) = 0$ for all b
- To avoid such problem, start with pseudocounts of r_{kl} for A_{kl} and $r_k(b)$ for $E_k(b)$.
- Large pseudocount indicates strong prior belief about the probabilities (will require more data to modify)
- Small pseudocount just to avoid zero probability



Saad Mehmeh

Example

Dishonest Casino HMM

$$r_{0F} = r_{0L} = r_{F0} = r_{L0} = 1; \quad [\text{avoid zero probability}]$$

$$r_{FL} = r_{LF} = r_{FF} = r_{LL} = 1; \quad [\text{avoid zero probability}]$$

$$r_F(1) = r_F(2) = \dots = r_F(6) = 20 \quad [\text{strong belief that fair is fair}]$$

$$r_L(1) = r_L(2) = \dots = r_L(6) = 5 \quad [\text{wait and see for loaded}]$$



Saad Mehmeh

New species comes in...

- New species with different distribution of CpG islands.
- We do not have labeled genomic sequences for the new species.
- Need to find maximum likelihood θ of HMM without knowing the paths!



Saad Mehmeh

Baum–Welsh algorithm

start at iteration 0 with some θ , call it θ^0

$$L^0 \leftarrow \log p(x^1 \dots x^n \mid \theta^0)$$

$i \leftarrow 0$

repeat

$i \leftarrow i + 1$

$$A_{kj}^i \leftarrow E[A_{kj} \mid x^1 \dots x^n, \theta^{i-1}] \quad (\text{expected value})$$

$$E_k(b)^i \leftarrow E[E_k(b) \mid x^1 \dots x^n, \theta^{i-1}] \quad (\text{expected value})$$

calculate θ^i using maximum likelihood estimators from counts A_{kj}^i and $E_k(b)^i$ as before.

$$L^i \leftarrow \log p(x^1 \dots x^n \mid \theta^i) \quad (\text{new likelihood})$$

until $L^i - L^{i-1} < \text{threshold}$



Saad Mehmeh

What is the guarantee?

- Baum–Welsh algorithm is a special case of a general algorithm known as Expectation Maximization (EM)
- EM guarantees that $p(X \mid \theta^{i+1}) \geq p(X \mid \theta^i)$
- It will therefore converge to a local maximum (not necessarily the maximum)



Saad Mehmeh

We need to...

Compute:

$$- E[A_{kl} | x^1 \dots x^n, \theta]$$

$$- E[E_k(b) | x^1 \dots x^n, \theta]$$



Saad Mehmeh

$$E[A_{kl} | x^1 \dots x^n, \theta]$$

By linearity of expectation:

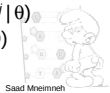
$$E[A_{kl} | x^1 \dots x^n, \theta] = \sum_j E[A_{kl}^j | x^j, \theta]$$

By linearity of expectation, again:

$$\begin{aligned} E[A_{kl}^j | x^j, \theta] &= \sum_i E[\# \text{ of } k \rightarrow l \text{ at } x_i^j | x^j, \theta] \\ &= \sum_i p(k \rightarrow l \text{ at } x_i^j | x^j, \theta) \end{aligned}$$

$$p(k \rightarrow l \text{ at } x_i^j | x^j, \theta) = p(\pi_i = k, \pi_{i+1} = l | x^j, \theta)$$

$$\begin{aligned} p(\pi_i = k, \pi_{i+1} = l | x^j, \theta) &= p(\pi_i = k, \pi_{i+1} = l, x^j | \theta) / p(x^j | \theta) \\ &= f_k^j(i) a_k e_l(x_{i+1}^j) b_l^j / (i+1) / p(x^j | \theta) \end{aligned}$$



Saad Mehmeh

We get:

$$E[A_{kl} | x^1 \dots x^j, \theta] = \sum_j \frac{1}{p(x^j | \theta)} \sum_i f_k^j(i) a_k e_l(x_{i+1}^j) b_l^j (i+1)$$

$$E[E_k(b) | x^1 \dots x^j, \theta] = \sum_j \frac{1}{p(x^j | \theta)} \sum_{i | x_i^j = b} f_k^j(i) b_k^j (i)$$



Saad Mehmeh

Viterbi training

start at iteration 0 with some θ , call it θ^0
 $i \leftarrow 0$

repeat
 $i \leftarrow i + 1$

$A_{kl}^i \leftarrow$ number of transitions $k \rightarrow l$ on the most probable paths $\pi^{1*}, \dots, \pi^{n*}$

$E_k(b)^i \leftarrow$ number of times k emits b on the most probable paths $\pi^{1*}, \dots, \pi^{n*}$

calculate θ^i using maximum likelihood estimators
from counts A_{kl}^i and $E_k(b)^i$ as before.

until none of the optimal paths change



Saad Mheineh

What is the guarantee

- It will converge
- It will not necessarily maximize the true likelihood $p(x_1 \dots x^n \mid \theta)$, but $p(x_1 \dots x^n \mid \theta, \pi^{1*}, \dots, \pi^{n*})$
- Usually performs less well than Baum–Welsh
- Practical, don't have to perform Forward and Backward algorithms, only Viterbi!
- Makes sense if we are using only Viterbi decoding



Saad Mheineh
