

CS 378 Lecture 18

Today

- Finish IBM Model 1
- Phrase-based MT
- Syntactic MT (briefly)
- Neural MT and seq2seq models

Recap IBM Model 1

$\bar{s} =$ Je fais un bureau NULL
 $\bar{a} =$ $\begin{array}{c} | a_1=1 \\ \swarrow a_2=2 \end{array}$ $\begin{array}{c} \searrow a_3=2 \\ \swarrow a_4=3 \end{array}$ $\begin{array}{c} \searrow a_5=4 \end{array}$
 $\bar{t} =$ I am making a desk

$$P(\bar{t}, \bar{a} | \bar{s}) = \prod_{i=1} P(a_i) P(t_i | s_{a_i})$$

$P(a_i) =$ uniform dist over $\downarrow$ NULL
 $\{1, 2, \dots, m+1\}$

$$P(t_i | s_{a_i}) \quad P(I | Je) = 0.8$$

...

Announcements

- Custom FP proposals returned w/ comments
- Optional lecture Tuesday
- A4
- A5

Inference in Model 1

What we care about: $P(\bar{a} | \bar{t}, \bar{s})$

$$\left(\underset{\bar{a}}{\operatorname{argmax}} P(\bar{a} | \bar{t}, \bar{s}) \right)$$

$$P(\bar{a} | \bar{t}, \bar{s}) = \frac{P(\bar{a}, \bar{t} | \bar{s})}{P(\bar{t} | \bar{s})} = \frac{1}{\prod_{i=1}^n P(a_i) P(t_i | s_i)} \underset{\text{constant}}{\quad}$$

$$P(\bar{a} | \bar{t}, \bar{s}) \text{ proportional to } \prod_{i=1}^n P(t_i | s_{a_i})$$

$$P(a_i | \bar{t}, \bar{s}) \text{ proportional to } P(t_i | s_{a_i})$$

(model params)

$P(I|S)$

I like eat

Je	0.8	0.1	0.1
Jl	0.8	0.1	0.1
mange	0	0	1.0
aine	0	1.0	0
NULL	0.4	0.3	0.4

Example

$\bar{s} = Je \text{ NULL}$

a_1

$\bar{t} = I$

$$P(a_i | \bar{t}, \bar{s}) \text{ prop. to } \begin{cases} a_1 = 1 & P(I|Je) = 0.8 \\ & Je \\ a_1 = 2 & P(I|NULL) = 0.4 \\ & NULL \end{cases}$$

$$P(a_i | \bar{t}, \bar{s}) = \begin{cases} a_1 = 1 & 2/3 \\ a_1 = 2 & 1/3 \end{cases}$$

Ex 2 J' aine NULL

I like

$$P(a_1 | \bar{s}, \bar{t}) = \begin{cases} 0.8 & J' & 2/3 \\ 0 & \text{aine} \Rightarrow 0 \\ 0.4 & \text{NULL} & 1/3 \end{cases}$$

$$P(a_2 | \bar{s}, \bar{t}) = \begin{cases} 0.1 & J' & 1/14 \\ 1.0 & \text{aine} \Rightarrow 10/14 \\ 0.3 & \text{NULL} & 3/14 \end{cases}$$

More complex models : HMM alignment model

$$P(a_i | a_{i-1})$$

IBM Models 2-4

Learning Unsupervised learning: $\bar{s}, \bar{t}$
no $\bar{a}$

EM (Expectation Maximization)

$$\text{maximize} \quad P(\bar{T}^{(i)} | \bar{S}^{(i)})$$

$$\sum_{i=1}^D \log \sum_{\bar{a}} P(\bar{a}, \bar{T}^{(i)} | \bar{S}^{(i)})$$

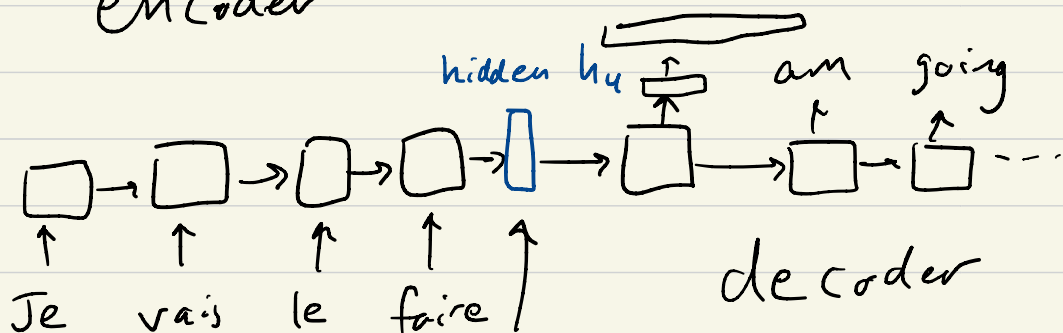
log marginal likelihood

SLIDES

Seq 2 seq models

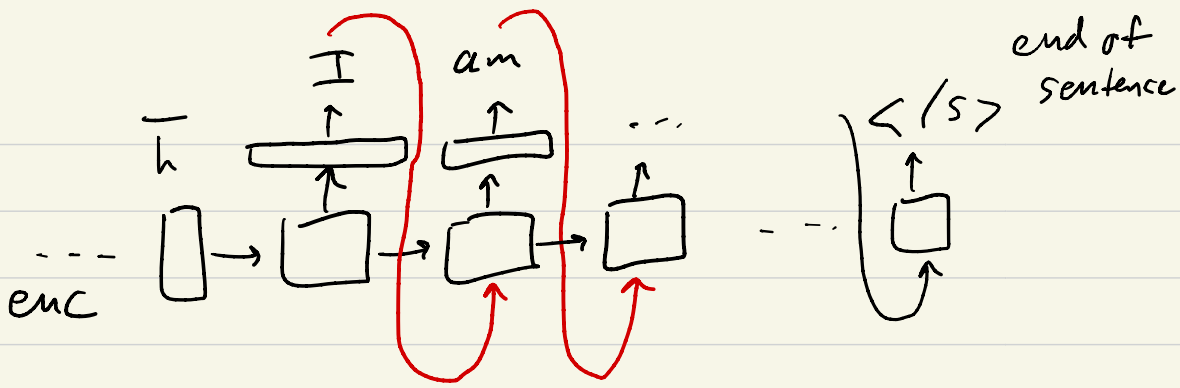
Key idea: encode source sent w/ RNN,
"decode" target w/ another RNN

encoder $I \leftarrow P(t_1 | \bar{s}) = \text{softmax}(w \bar{h}_1)$



representation of the French
sentence

$$P(\bar{T} | \bar{s}) = P(t_1 | \bar{s}) P(t_2 | \bar{s}, t_1) \dots$$



Feed output t_i into cell input for t_{i+1}

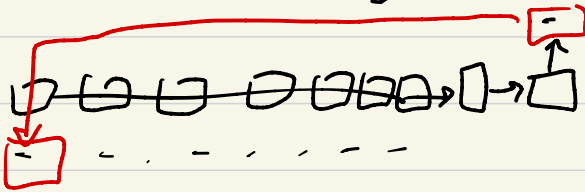
Why? Model capture word-word dependencies more easily

Training Given sequence pairs $(\bar{s}, \bar{t})$

$$\text{loss} = \sum \log P(t_i^* | \bar{s}, t_{1, \dots, i-1}^*)$$

Like in language modeling, assume everything up until now matches our reference / gold

Problem Long-range dependencies



What we want is a way to
look back at the input more easily

Solution: attention