

Announcements

- A4 due tonight
- A5 out tonight, due next Tuesday

Recap Phrase-based MT

- ① Word alignment ($src \rightarrow trg / trg \rightarrow src$, intersect)
- ② Phrase extraction
- ③ Decoding: search over phrase lattice + LM

Today Seq-to-seq models / Neural MT

↳ Connections to RNN LMs

↳ Training

↳ Problems

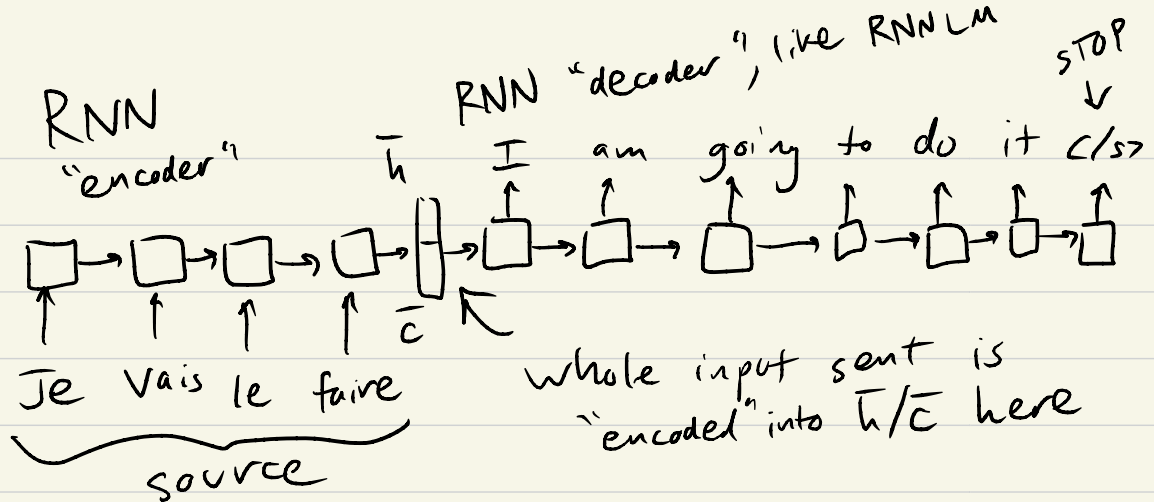
Attention

↳ Mechanism

↳ Connections to alignment

Seq-to-seq models for MT

Sutskever et al. NeurIPS 2014



NMT: encode src w/RNN, "decode" trg with another RNN (diff params)

$$P(\bar{T} | \bar{S}) = P(t_1 | \bar{S}) P(t_2 | t_1, \bar{S}) P(t_3 | t_{1:2}, \bar{S}) \dots$$

$$P(t_i | t_{1:i-1}, \bar{S}) = \text{Softmax} (W^{(d)} \bar{h}_i^{(d)})$$

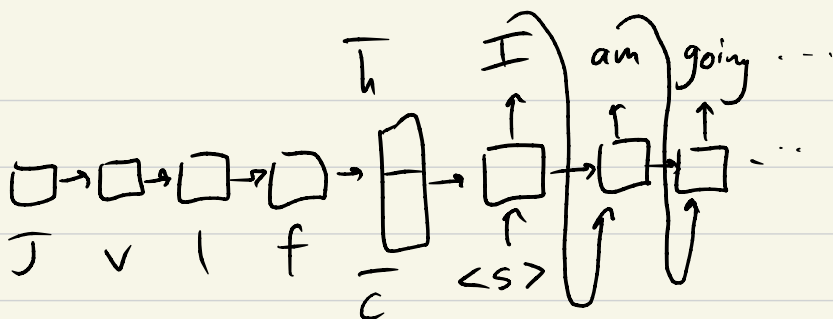
$W^{(d)}$: trg vocab

x hidden size mat

hidden state at i th step of decoder

Like an RNN LM additionally conditioned on $\bar{S}$

Decoder: feed the previous predicted word into next decoder timestep

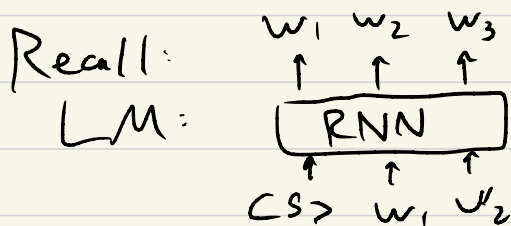


Decoding: $t_0 = \langle s \rangle$

while $t_i \neq \text{STOP}$

$$t_i = \arg \max_t P(t | t_{1:i-1}, \bar{s})$$

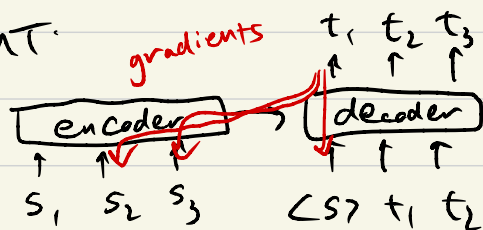
Training Looks like LM training



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

t_i^* gold
 t_i trg

NMT:



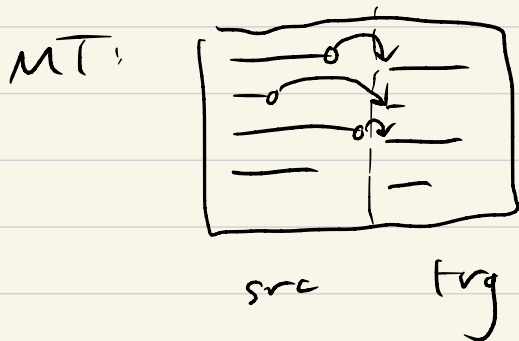
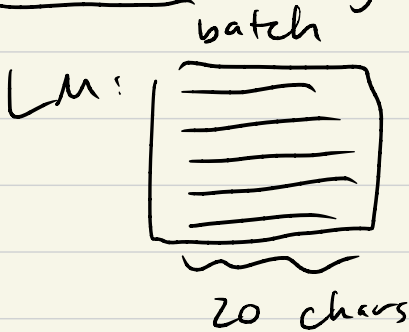
encoder trained too!

teacher forcing:
feed in true
target tokens

Teaching the model to predict t_i | the correct t_{i+1} up to this point

What do we really want? Model's predicted sequence $\hat{t}_1 \dots \hat{t}_n$ has high BLEU score
Requires reinforcement learning

Details Lengths can vary, need to pad



need padding + smart handling

What goes wrong?

① Repetition

Je vais le faire → I am going to
do it going to going to
going to ...

Won't happen in PBMT (every word
translated as part of one phrase)
Need a notion of coverage in NMT

② Unknown words

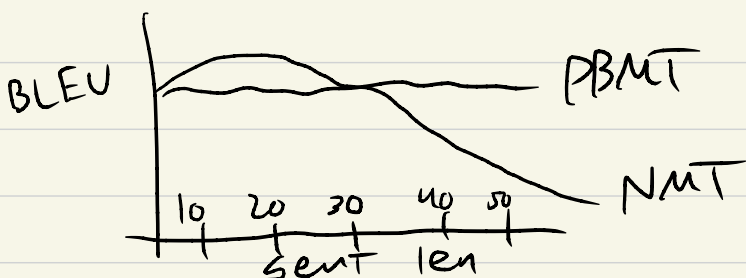
Decoder has a fixed vocab $|V_t|$

PBMT: "copy" rules

Pont-de-Buis → copy to output

NMT: produce UNK

- ① Repetition (going to going to)
- ② UNKS
- ③ Poor performance on long sent



Bahdanau et al. (2015)
(attention)

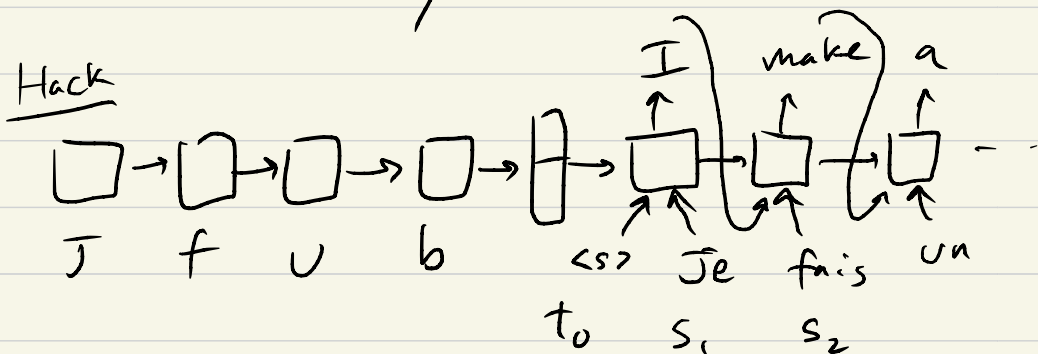
NMT: fixed size $\vec{h}$ / $\vec{c}$ vectors don't scale arbitrarily
LSTMS can't remember for too long

Sutskever: reverse input as a trick to address this

$s_m, s_{m-1}, \dots, s_1 \rightarrow t_1, t_2, \dots$

at least remember the first few $\vec{s}$

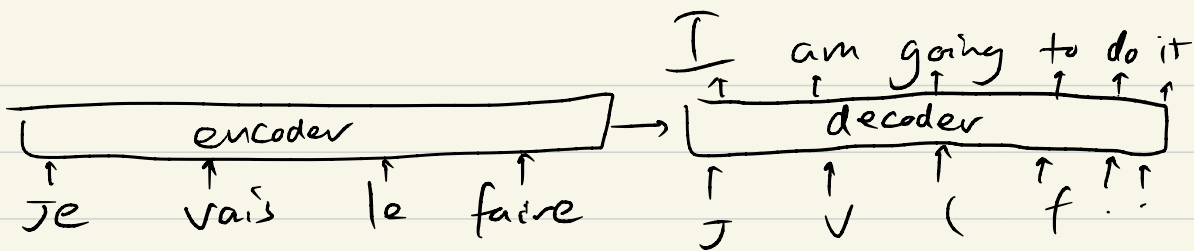
Attention: Je fais un bureau → I make a desk
 What if we know translation should be word-by-word?



Manually feed s_i into i th decoder step
 (along with t_{i-1})

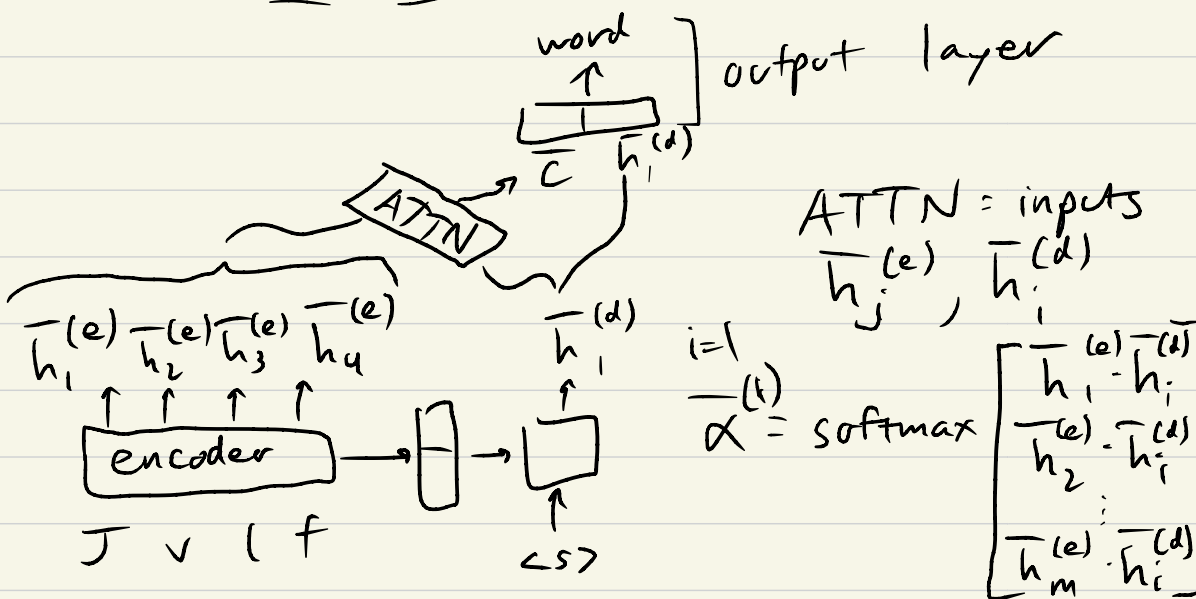
Q: What problems can this fix? (1-3)

- (+3) ✓ ① Less likely to repeat if we see explicit words
- ✗ ② Requires changing output layer
- ✓ ③ Anchored more to input

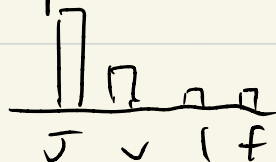


- Run out of s before end of t ?
- Ordering is inflexible

Attention Target decoder picks where to look in source



α distribution over source posns



$$\bar{c}_i = \sum_{j=1}^m \overset{\text{scalar}}{\alpha_j^{(i)}} \cdot \overset{\text{vector}}{\bar{h}_j^{(e)}}$$

weighted average
of $\bar{h}^{(e)}$ with $\bar{\alpha}$
as weights

$$P(t_i | \bar{s}, t_{1:i-1}) = \text{softmax} \left(W^{(d)} \begin{bmatrix} \bar{h}_i^{(d)} \\ \bar{c}_i \end{bmatrix} \right)$$

Concat

$\bar{\alpha}$ can also use other fns
to compare $\bar{h}^{(d)} / \bar{h}^{(e)}$ besides dot

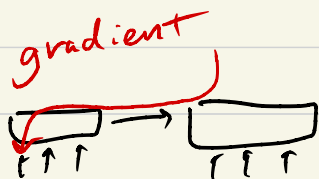
Think about Je fais un bureau

$$\alpha_1 = [0.99, 0.003, 0.003, 0.003]$$

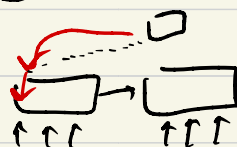
Output: $\text{softmax} \left(W^{(d)} \begin{bmatrix} \bar{h}_1^{(d)} \\ \approx \bar{h}_1^{(e)} \end{bmatrix} \right)$

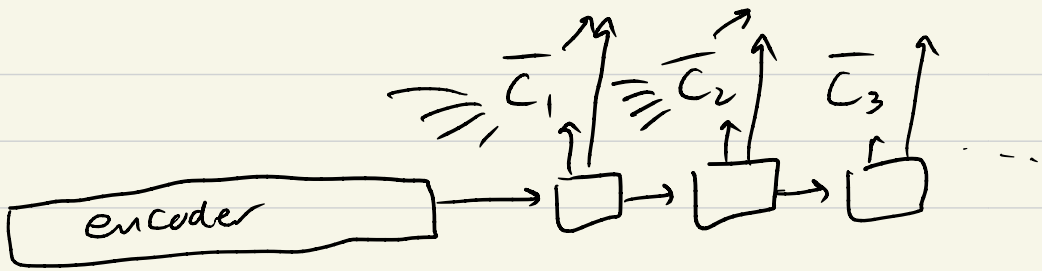
$\approx$ Je

No attn:



w/attn





Je vais le faire

I am going to do it

Learning mapping 1 2 2 4 4 3

- Not input at train time!

Je vais le faire → 1 2 2 4 4 3

- easier than translation
- regular mapping
- This is learnable!