

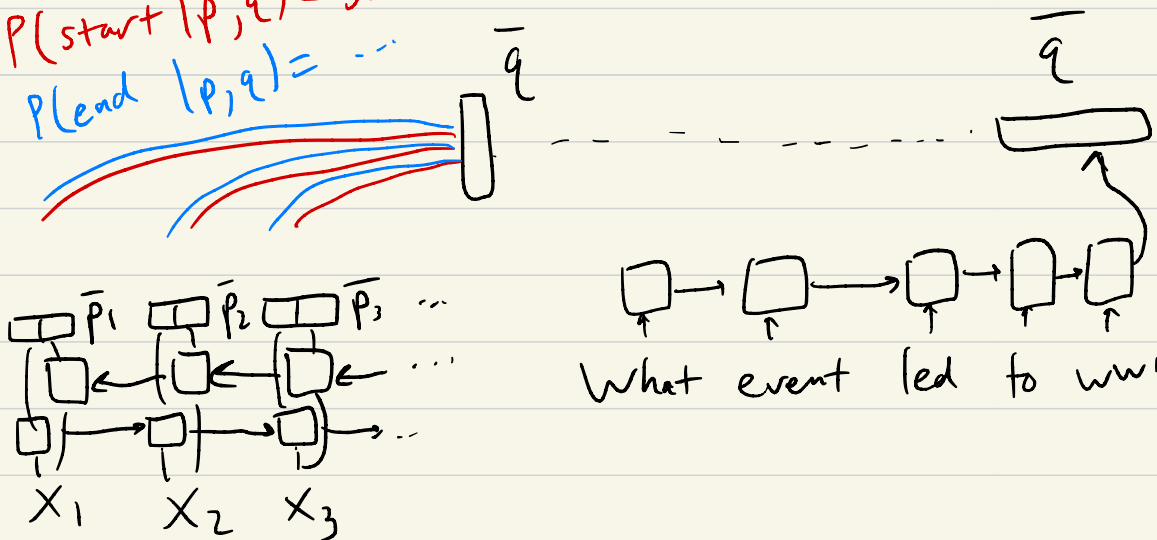
Announcements

- A4, A5 grading underway
- FP check-in May 1

Recap Reading comprehension: Attentive reader

$$P(\text{start} | p, q) = \text{softmax}(\bar{p}_i w^{\text{start}} \frac{q}{q})$$

$$P(\text{end} | p, q) = \dots$$



The a. of F F caused ..

Input: doc and question

Output: distributions over start and end points of answer

- Today
- ① Two improvements to Attentive Reader
 - ② SQuAD state-of-the-art
 - ③ Pre-training with ELMo

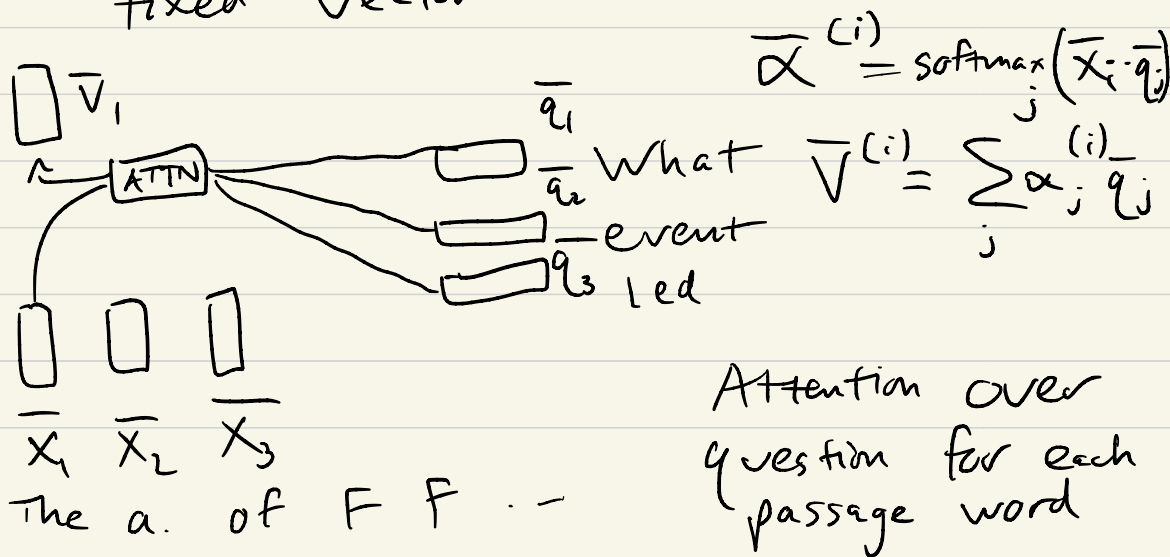
Improvements to AR

$\bar{p}, \bar{q}$ encoders more complex

① "context to query" attention

Intuition: reducing the question to $\bar{q}$ throws out info

→ hard to represent entities in a fixed vector



indices $1 \dots m$

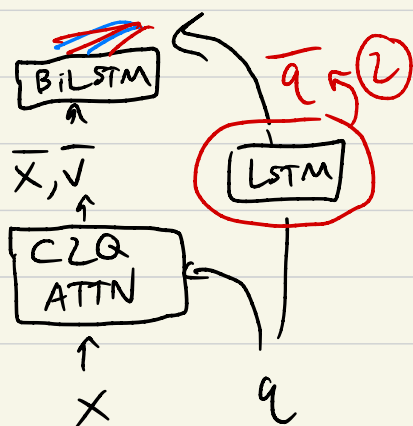
$\bar{x}^{(i)}$: distribution over $\bar{q}_1 \dots \bar{q}_m$
for the i th passage token

$\bar{v}^{(i)}$: representation of the question
"tailored" to this passage token

Suppose $\bar{x}_8 = WW1$
 $\bar{q}_7 = WW1$

$\bar{v}^{(8)} \approx WW1 \approx \bar{x}^{(8)}$

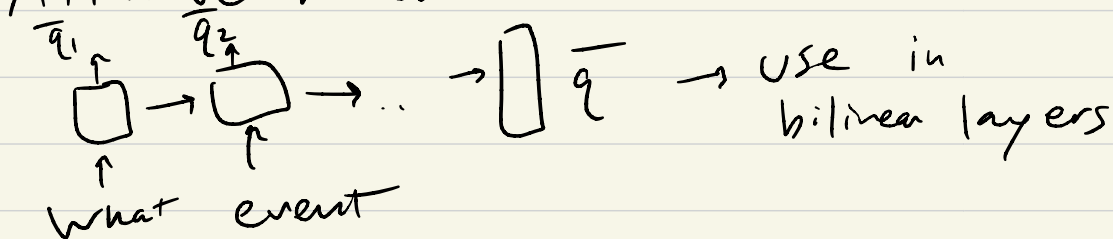
We can see at this point that the
question has matched the passage



By injecting info about
 $\bar{q}$ earlier, the model
can do better at
finding the right spots

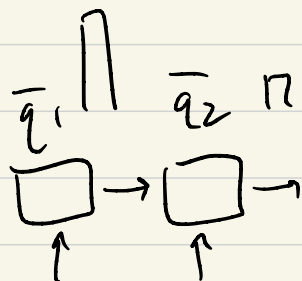
② question encoding modification

Attentive reader



Switch to a more "bag of words" like representation

$$b_j = \text{softmax}_j \left(\underbrace{\bar{w}}_{\text{trainable param vector}} \cdot \bar{q}_j \right) \quad \begin{array}{l} \text{dist over} \\ q \text{ tokens} \end{array}$$



what event ..

$$\bar{q} = \sum b_j \bar{q}_j$$

looks like attn
but uses a fixed
"key" $\bar{w}$

"span attention"

