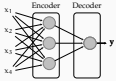


- All machine learning methods can be seen performing two tasks:
 - encoding the input into a useful representation
 - decoding the representation into the output
- In more traditional methods, encoding is 'manual', or external to the learning algorithm
- Modern deep learning methods include the encoder: they learn to build (multiple layers/hierarchies of) useful input representations

A simple encoder-decoder network

we can view any neural network as two parts

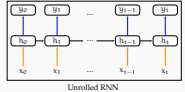


- The encoder encodes the input in hidden representations
- Decoder 'decodes' the encoded input to the output
- In computational linguistics, many tasks require encoding and decoding sequences

Recap: RNNs



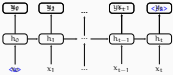
RNN



Unrolled RNN

Uses of RNNs

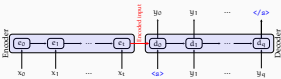
Many-to-many with a delay (e.g., machine translation, summarization, question answering, ...)



RNN problems, solutions, variations

- (Unfolded) RNNs can be deep (based on the length of the input), this results in
 - exploding gradients – solution: gradient clipping
 - vanishing gradients – solution: gated RNNs (to some extent)
- More generally, keeping relevant information over longer sequences are difficult – solution: attention (this lecture)
- RNNs condition the prediction in only one direction – solution: bidirectional RNNs
- It is also common to stack multiple layers of RNNs
- RNNs are inherently sequential, this prevents parallel processing

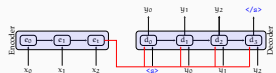
Sequence-to-sequence RNNs (seq2seq)



- Note that the decoder is a RNN language model
- Both input and output can be arbitrary length
- All information about the input is coded in a single encoding vector

Sequence-to-sequence RNNs (seq2seq)

a simple improvement

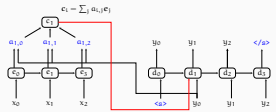


- Instead of passing the encoder output (context vector) to the first decoder state, pass it to all time steps
- Helps decoder to not to forget the encoder state, but early words in the encoder may not be represented well in the context vector

Attention: the general idea

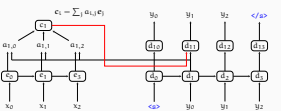
- A naive solution could be passing all intermediate steps of the encoder to the decoder states (e.g., concatenate or average)
- However, in many problems, a part of the input is relevant for predicting the current output
- A common solution to the problem is *attention* mechanism
- Attention is about focusing into the relevant parts of the input
- In sequence-to-sequence problems, this corresponds to alignment

Sequence-to-sequence RNN with attention



Sequence-to-sequence RNN with attention

another variation



Calculating attention weights

- The context vector is the sum of the encoder states, weighted by attention, $a_{1,j}$

$$c_1 = \sum_j a_{1,j} e_j$$

- Typically the weights are normalized through *softmax*: the result is an attention distribution

$$a_{1,j} = \frac{e^{f(d_{1,j}, e_j)}}{\sum_k e^{f(d_{1,k}, e_k)}}$$

- The attention function, $f(\cdot)$ computes the relevance of encoder state e_j to the decoder state d_1

Some (common) attention functions

- Dot product:

$$f(\mathbf{d}_i, \mathbf{e}_j) = \mathbf{d}_i^T \mathbf{e}_j$$

- Generalized dot product:

$$f(\mathbf{d}_i, \mathbf{e}_j) = \mathbf{d}_i^T \mathbf{W}_a \mathbf{e}_j$$

- Scaled dot product:

$$f(\mathbf{d}_i, \mathbf{e}_j) = \frac{\mathbf{d}_i^T \mathbf{e}_j}{\sqrt{k}}$$

- Additive attention:

$$f(\mathbf{d}_i, \mathbf{e}_j) = v^T \tanh(\mathbf{W}_a \mathbf{d}_i + \mathbf{U}_a \mathbf{e}_j)$$

Attention and content addressable memory

- In the literature, attention mechanism is often explained as a form of content-addressable (or associative) memory: decoder state is used to query the encoder states

$$a_{i,j} = \frac{e^{f(\mathbf{d}_{i-1}, \mathbf{e}_j)}}{\sum_k e^{f(\mathbf{d}_{i-1}, \mathbf{e}_k)}} \quad c_i = \sum_j a_{i,j} \mathbf{e}_j$$

- In this setting, *key* and *value* are the same
- In case of hard attention, the mechanism is equivalent to associative arrays (maps)

Example attention weights

caption generation



Summary

- Attention is a general mechanism to focus on certain parts of input
 - It is also useful for generating 'explanations'
 - Attention is the basic mechanism behind the current state of the art models (transformers)
 - Reading: Jurafsky and Martin (2025, Chapter 8)
- Next:
- Self attention and transformers architecture (Reading: Jurafsky and Martin (2025, Chapter 9))

Additional reading, references, credits (cont.)

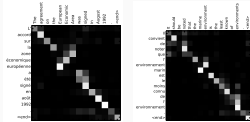
- Nishitane, Shinya, Kyunghyun Cho, and Yoshua Bengio (2014). "Neural machine translation by jointly learning to align and translate". In: *arXiv preprint arXiv:1409.0329*.
- Vaswani, Ashish, and James H. Martin (2020). *Speech and Language Processing: The Introduction to Natural Language Processing, Computational Linguistics, and Speech Recognition with Language Models, 3rd Edition*. Morgan Kaufmann, 2020.
- Yu, Kehua, Junyi He, Ryan Kiros, Kyunghyun Cho, Aaron Courville, Radu S. Susskind, Rick J. Zemel, and Yoshua Bengio (2015). "Bridging visual and text neural image caption generation with visual attention". In: *International conference on machine learning*. PMLR, pp. 2240-2247.

Hard and soft attention

- The mechanism we described is called *soft attention*: The attention mechanism allows attending to more than one input in a weighted manner
- The case where the attention weights form a one-hot vector is called *hard attention*
- Soft attention is common, both because of its flexibility and ease of training (continuous / differentiable functions)

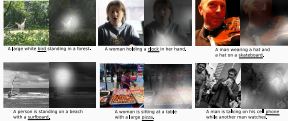
Example attention weights

machine translation (en-fr)



Example attention weights

caption generation



Additional reading, references, credits

- The translation example is from Bahdanau, Cho, and Bengio (2014)
- The image captioning examples are from Xu et al. (2015)