

Sequence-to-sequence (encoder-decoder) networks

Statistical Methods in NLP 2

ISCL-BA-08

Çağrı Çöltekin
ccoltekin@fsa.uni-tuebingen.de

University of Tübingen
Seminar für Sprachwissenschaft

Summer Semester 2026

version: 2024-08-19 10:00:00 (UTC)

Encoder-decoder models

- 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

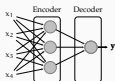
Ç. Çöltekin, SS | University of Tübingen

Summer Semester 2026 1 / 18

Recap: RNNs Encoder-decoder RNNs Attention Summary

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

Ç. Çöltekin, SS | University of Tübingen

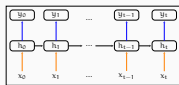
Summer Semester 2026 2 / 18

Recap: RNNs Encoder-decoder RNNs Attention Summary

Recap: RNNs



RNN



Unrolled RNN

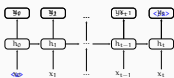
Ç. Çöltekin, SS | University of Tübingen

Summer Semester 2026 3 / 18

Recap: RNNs Encoder-decoder RNNs Attention Summary

Uses of RNNs

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



Ç. Çöltekin, SS | University of Tübingen

Summer Semester 2026 4 / 18

Recap: RNNs Encoder-decoder RNNs Attention Summary

RNN problems, solutions, variations

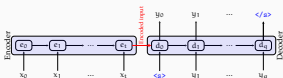
- (Unfolded) RNNs can be deep (based on the length of the input), this results in
 - exploding gradients – solution: gradient clipping
 - vanishing gradients – solution: gated RNNs (to some extent)
- More generally, keeping relevant information over longer sequences are difficult – solution: attention (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

Ç. Çöltekin, SS | University of Tübingen

Summer Semester 2026 5 / 18

Recap: RNNs Encoder-decoder RNNs Attention Summary

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

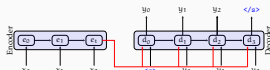
Ç. Çöltekin, SS | University of Tübingen

Summer Semester 2026 6 / 18

Recap: RNNs Encoder-decoder RNNs Attention Summary

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

Ç. Çöltekin, SS | University of Tübingen

Summer Semester 2026 7 / 18

Recap: RNNs Encoder-decoder RNNs Attention Summary

Attention: the general idea

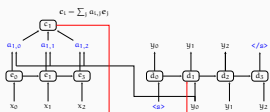
- A naive solution could be passing all intermediate steps of the encoder to the decoder states (e.g., concatenate or average)
- However, in many problems, a part of the input is relevant for predicting the current output
- A common solution to the problem is *attention* mechanism
- Attention is about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

Ç. Çöltekin, SS | University of Tübingen

Summer Semester 2026 8 / 18

Recap: RNNs Encoder-decoder RNNs Attention Summary

Sequence-to-sequence RNN with attention



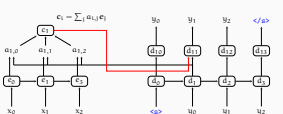
Ç. Çöltekin, SS | University of Tübingen

Summer Semester 2026 9 / 18

Recap: RNNs Encoder-decoder RNNs Attention Summary

Sequence-to-sequence RNN with attention

another variation



Ç. Çöltekin, SS | University of Tübingen

Summer Semester 2026 10 / 18

Recap: RNNs Encoder-decoder RNNs Attention Summary

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

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

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

Ç. Çöltekin, SS | University of Tübingen

Summer Semester 2026 11 / 18

