



AI and Deep Learning

5-1 Recurrent Neural Network I

Zhonglei Wang

WISE, SOE, CHOW

XMU

(Version v1)

Contents

1. Motivation

2. Notation

3. Recurrent neural networks

Motivation

1. It is a common task for natural language processing (NLP)
 - Machine translation
 - Sentiment analysis
 - ChatGPT (Chat Generative Pre-trained Transformer)
2. For example, find the destination in the following sentences
 - I arrive at **Beijing** from Xiamen
 - I leave Beijing to **Xiamen**
3. Problems
 - How to convert sentences into features for calculation?
 - What are appropriate models for NLP?

Motivation

1. We usually use the following steps for NLP

- Tokenization (For feature extraction)
- Embedding
- Modeling (For analysis)

Tokenization

1. For English, tokenization is straightforward
2. For example, consider “I arrive at Beijing from Xiamen.”
3. We can tokenize the above sentence into the following **seven** parts

I / arrive / at / Beijing / from / Xiamen / .

4. Difficulties for **Tokenization**

- Some words, including names or rarely used ones, may not be in the vocabulary
- It is not clear how to handle punctuations, including “.,;:?! ”
- Different tokens may be needed for the same word with different suffixes
 - ▷ For example, walk, walks, walked...

Tokenization

1. **Possible solution:** A vocabulary includes
 - Commonly used words
 - Word fragments from which larger or less frequent words can be formed
2. Then, a vector is assigned to each token according to a vocabulary

One-hot encoding

1. In practice, we have a vocabulary of size $N(\approx 30,000)$
2. One-hot encoding can be used to represent each token in this vocabulary
 - A vector of length N
 - Contains 0 but only a single 1, indicating the position of that token in the vocabulary
 - The index of 1 shows the position of that token in the vocabulary

Remarks

1. Disadvantages

- **High dimensionality and sparsity**: Most information is redundant, especially when the vocabulary is large
- **Overfitting** due to the large dimension of the one-hot encoding
- **Cannot model relationship** between (among) words, such as “Man-Woman”
- **Sensitive to small changes** in the vocabulary: adding or deleting a word may change all encoding system
-

Embedding

1. We want a new token representation (**Embedding**)

- Achieving low dimensionality (≈ 1024)
- Reflecting relationship between (among) words

2. Possible solutions

- Word2Vec: CBOW+Skip-gram
- GloVe: Generalizes Word2Vec
- N-Gram: A probability model
- TF-IDF
- BERT

Embedding

1. Embedding is a fundamental task for NLP
 - Embeddings are learned from data by models
 - For example, with embeddings, OpenAI's GPT models can generate more **coherent** and **contextually relevant** responses to user prompts and questions
2. Nevertheless, we may not discuss those techniques, and check by yourselves
3. In the following analysis, we assume that embedding is done for each word

Model

1. Find the destination of the following sentence

I arrive at **Beijing** from Xiamen

Embedding:

$\mathbf{x}^{<1>}$

$\mathbf{x}^{<2>}$

$\mathbf{x}^{<3>}$

$\mathbf{x}^{<4>}$

$\mathbf{x}^{<5>}$

$\mathbf{x}^{<6>}$

Label:

$y^{<1>} = 0$

$y^{<2>} = 0$

$y^{<3>} = 0$

$y^{<4>} = \mathbf{1}$

$y^{<5>} = 0$

$y^{<6>} = 0$

2. What is an appropriate model for NLP?

Model

1. Why not using FNN or CNN?

- Inputs may have different length
- Specifically, standard NN cannot deal with dependence

2. A desired model should satisfy

- It can handle sequences with different lengths
- It keeps information from context
- It generates different forms of outputs based on different learning goals

Recurrent neural network

1. A similar concept was proposed in 1980s
2. We focus on finding the destination in a sentence
 - Essentially, this is a binary classification task for each word in a sentence
3. We have a training dataset, consisting of different labeled sentences
4. Different tasks are discussed later

Recurrent neural network

1. Suppose we have a sentence

I arrive at Beijing from Xiamen

Embedding:

$\mathbf{x}^{<1>}$

$\mathbf{x}^{<2>}$

$\mathbf{x}^{<3>}$

$\mathbf{x}^{<4>}$

$\mathbf{x}^{<5>}$

$\mathbf{x}^{<6>}$

Label:

$y^{<1>} = 0$

$y^{<2>} = 0$

$y^{<3>} = 0$

$y^{<4>} = 1$

$y^{<5>} = 0$

$y^{<6>} = 0$

2. We are interested in estimating probabilities $\{\hat{y}^{<i>} : i = 1, \dots, 6\}$ for each word in this sentence

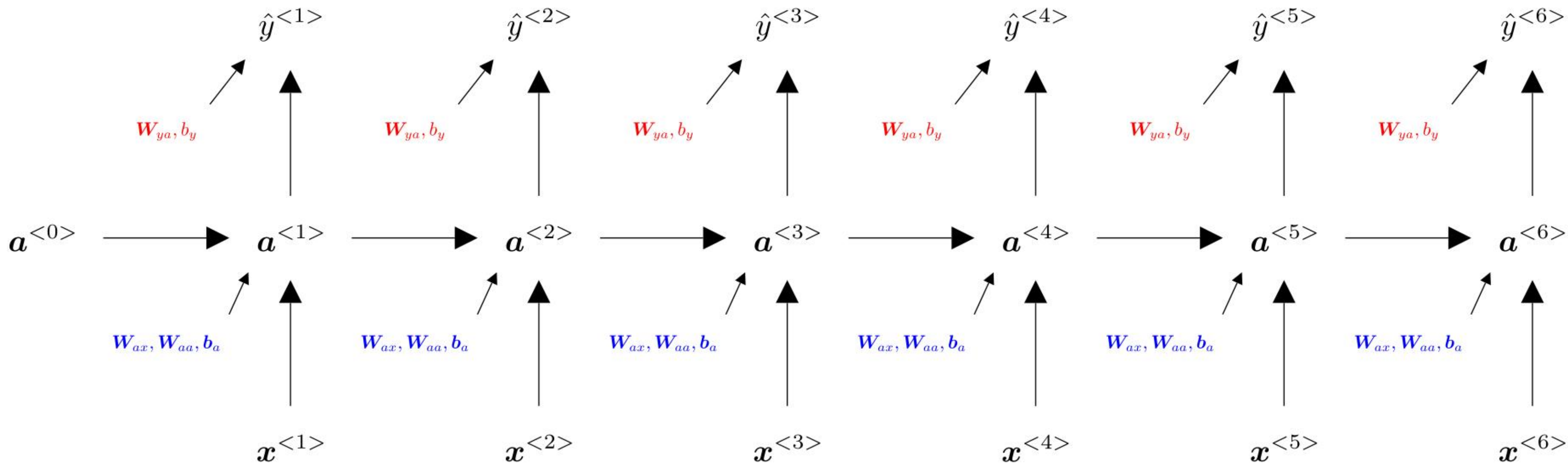
Building block

1. Initialize $\mathbf{a}^{<0>} = 0$
2. Given model parameters $\{\mathbf{W}_{ax}, \mathbf{W}_{aa}, \mathbf{b}_a, \mathbf{W}_{ya}, \mathbf{b}_y\}$, obtain

$$\begin{aligned}\mathbf{a}^{<i>} &= \sigma_{ax}(\mathbf{W}_{ax} \cdot \mathbf{x}^{<i>} + \mathbf{W}_{aa} \cdot \mathbf{a}^{<i-1>} + \mathbf{b}_a) \quad (i = 1, \dots, 6), \\ \hat{y}^{<i>} &= \sigma_{ya}(\mathbf{W}_{ya} \cdot \mathbf{a}^{<i>} + \mathbf{b}_y) \quad (i = 1, \dots, 6)\end{aligned}$$

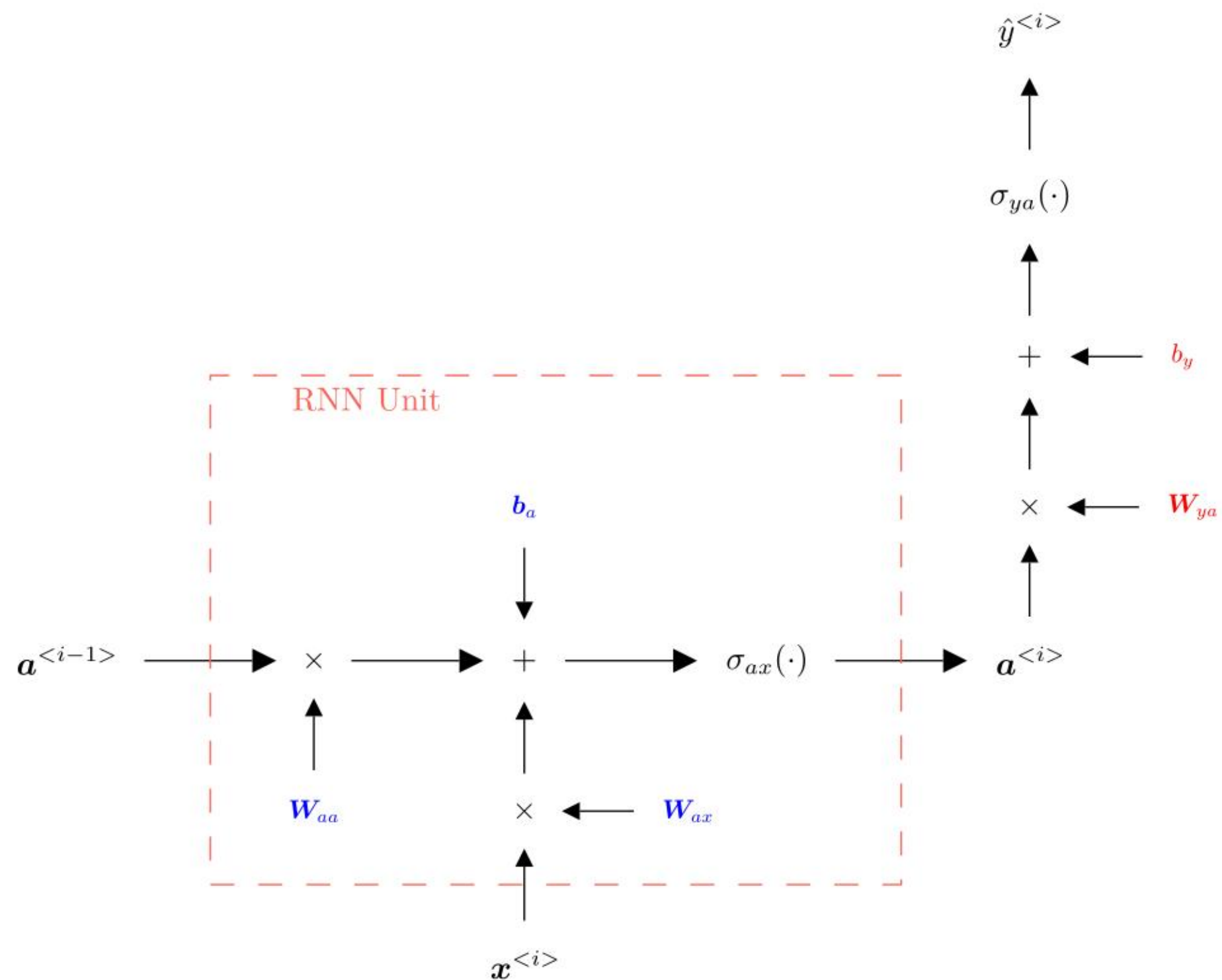
Building block

$$\hat{y}^{<i>} = \sigma_{ya}(\mathbf{W}_{ya} \cdot \mathbf{a}^{<i>} + \mathbf{b}_y) \quad (i = 1, \dots, 6)$$



$$\mathbf{a}^{<i>} = \sigma_{ax}(\mathbf{W}_{ax} \cdot \mathbf{x}^{<i>} + \mathbf{W}_{aa} \cdot \mathbf{a}^{<i-1>} + \mathbf{b}_a) \quad (i = 1, \dots, 6)$$

Flowchart



Remarks

1. Since we handle a binary classification problem, cross entropy is used as the loss function
2. Backpropagation can be derived based on the forward propagation in the previous slide
 - Remember: Sum up all derivatives involving information about the model parameter

Other examples

1. Sentiment analysis (Many to one)

- Feature: A sentence like “I like this movie very much”
- Label: Score like “5”

2. Translation (many to many)

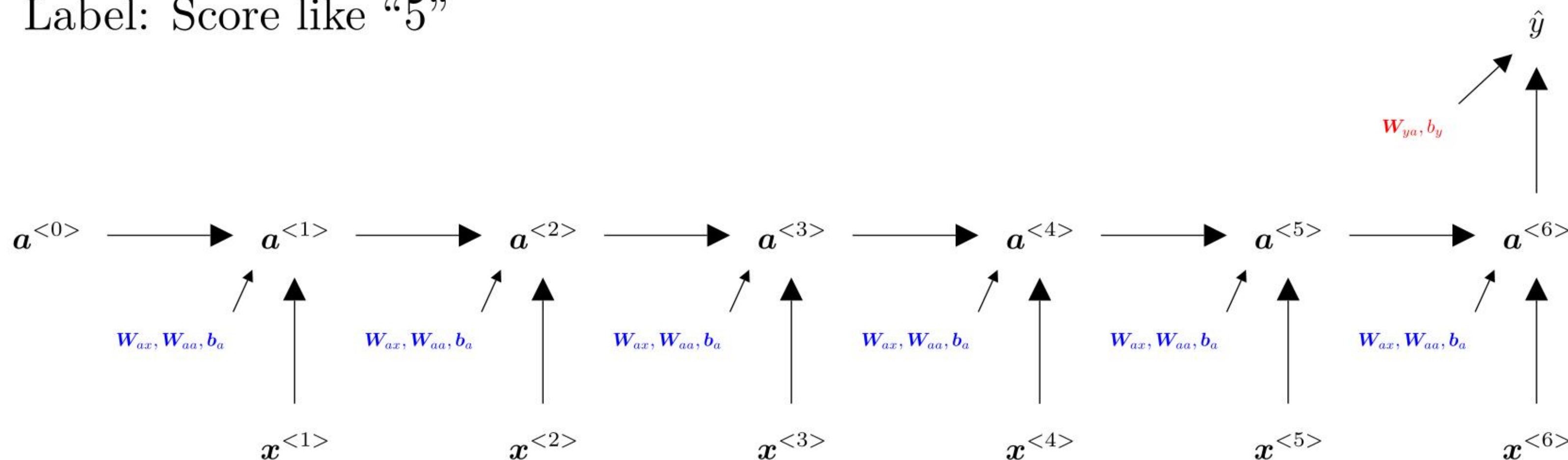
- Feature: A sentence
- Label: Translated sentence

3. Text generation (? to many)

- Feature: A start of a sentence like “I like ”
- Label (finish this sentence)

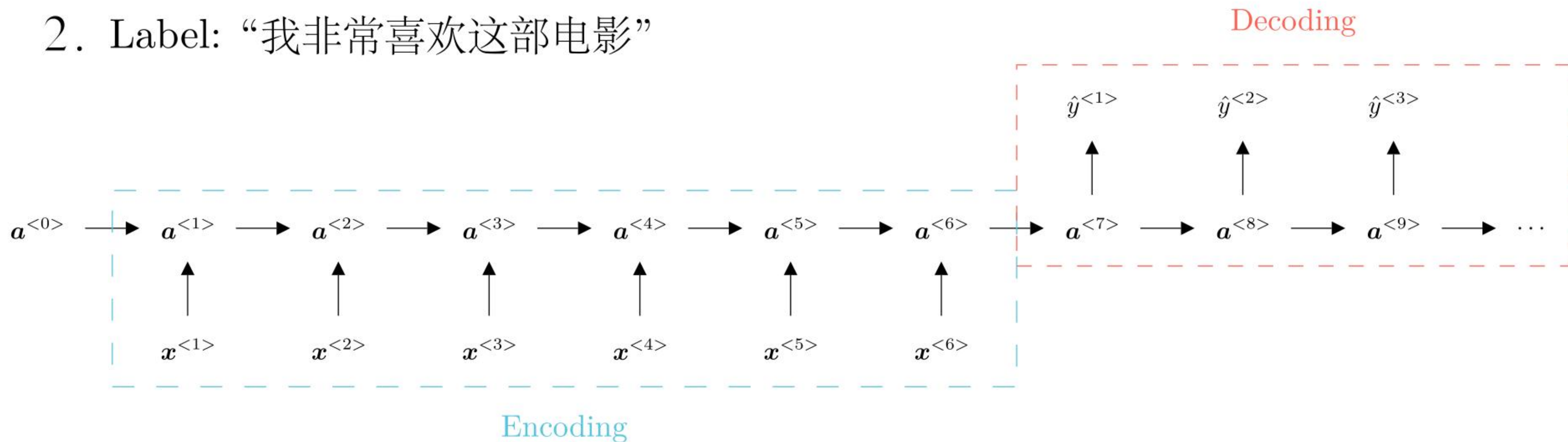
Sentiment analysis

1. Feature: A sentence like “I like this movie very much”
2. Label: Score like “5”



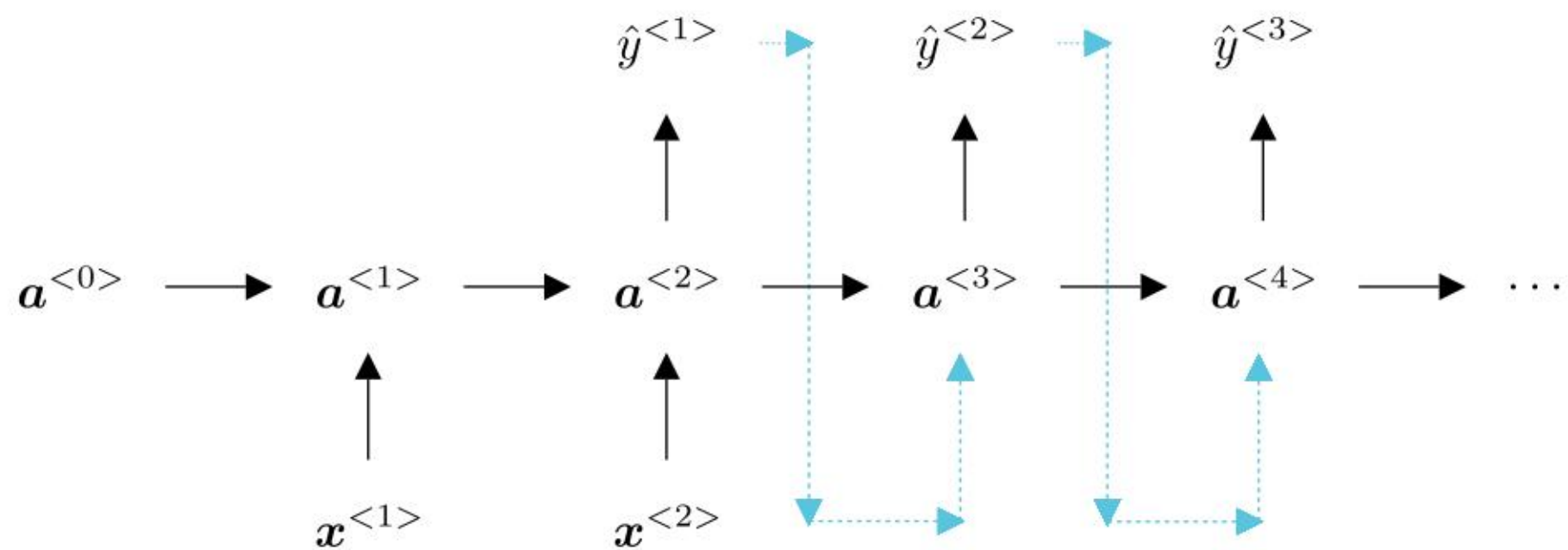
Translation

1. Feature: A sentence like “I like this movie very much”
2. Label: “我非常喜欢这部电影”



Text generation

1. Feature: A start of a sentence like “I like ”
2. Label (finish this sentence)



Remarks

1. Bengio et al. (1994) pointed out that an **RNN cannot capture long-term dependencies** since the gradients may either vanish (most of the time) or explode (rarely, but with severe effects).
2. “This makes gradient-based optimization method struggle, not just because of the variations in gradient magnitudes but because of the effect of **long-term dependencies** is hidden (being exponentially smaller with respect to sequence length) by the effect of **short-term dependencies**” (Chung et al., 2014)

Summary

1. Natural language is a sequential data
2. High-dimensional sparse vector representations is not suitable for NLP
3. RNN can be used to analyze sequential data
4. What are drawbacks of RNN