

Sentiment Classification Using Tree-Based Gated Recurrent Units

Vasileios Tsakalos

Dissertation presented as partial requirement for obtaining
the Master's degree in Information Management

Sentiment Classification Using Tree-Based Gated Recurrent Units

Copyright © Vasileios Tsakalos, NOVA Information Management School, NOVA University of Lisbon.

The NOVA Information Management School and the NOVA University of Lisbon have the right, perpetual and without geographical boundaries, to file and publish this dissertation through printed copies reproduced on paper or on digital form, or by any other means known or that may be invented, and to disseminate through scientific repositories and admit its copying and distribution for non-commercial, educational or research purposes, as long as credit is given to the author and editor.

ACKNOWLEDGEMENTS

I would like to thank my supervisor professor, Dr. Roberto Henriques, for helping me structure the report, advising me with regards to the mathematical complexity required. Moreover I would like to thank my brother, Evangelos Tsakalos, for providing me with the computational power to run the experiments and advising me regarding the subject of my master thesis.

ABSTRACT

Natural Language Processing is one of the most challenging fields of Artificial Intelligence. The past 10 years, this field has witnessed a fascinating progress due to Deep Learning. Despite that, we haven't achieved to build an architecture of models that can understand natural language as humans do. Many architectures have been proposed, each of them having its own strengths and weaknesses. In this report, we will cover the tree based architectures and in particular we will propose a different tree based architecture that is very similar to the Tree-Based LSTM, proposed by Tai(2015). In this work, we aim to make a critical comparison between the proposed architecture -**Tree-Based GRU**- with Tree-based LSTM for sentiment classification tasks, both binary and fine-grained.

Keywords: Deep Learning, Natural Language Processing, Recursive Neural Networks, Sentiment Classification

RESUMO

CONTENTS

List of Figures	xiii
List of Tables	xv
1 Introduction	1
1.1 Motivation	1
1.2 Overview of Thesis	1
1.3 Background	2
1.3.1 Sentiment Classification and Sequences	2
1.3.2 Syntactic Structure	2
1.3.3 Neural Networks	3
2 Word Embeddings	7
2.1 Latent Semantic Analysis	8
2.2 Word2Vec	9
2.3 GloVe	11
3 Feedback Neural Networks	13
3.1 Recurrent Neural Networks	13
3.1.1 Simple Recurrent Neural Networks	13
3.1.2 Long-Short Term Memory	15
3.1.3 Gated Recurrent Unit	17
3.2 Recursive Neural Networks	18
3.2.1 Simple Recursive Neural Network	19
3.2.2 Syntactically Untied SU-RNN	20
3.2.3 Matrix-Vector Recursive Neural Networks	21
3.2.4 Recursive Neural Tensor Network	23
3.2.5 Tree-Based Long-Short Term Memory Networks	24
3.2.6 Tree-Based Gated Recurrent Unit	27
4 Neural Network Training	29
4.1 Gradient Descent Variants	29
4.1.1 Batch Gradient Descent	30

4.1.2	Stochastic Gradient Descent	30
4.1.3	Mini-Batch Gradient Descent	30
4.2	Gradient Descent Extensions	31
4.2.1	Vanilla update	31
4.2.2	Momentum update	31
4.2.3	Nesterov Momentum update	31
4.2.4	AdaGrad	32
4.2.5	AdaDelta	33
4.2.6	RMSprop	34
4.3	Hyperparameters	34
4.3.1	Learning Rate	34
4.3.2	Regularization	34
5	Experiments	37
5.1	Model comparison	37
5.1.1	Classification model	37
5.1.2	Binary Classification	38
5.1.3	Fine-grained Classification	39
5.1.4	Hyperparameters and Training Details	39
5.2	Results	40
5.3	Conclusions and Future Work	41
	Bibliography	43
A	Appendix 1	51
B	Appendix 2	53

LIST OF FIGURES

1.1	Constituency Tree	3
1.2	Dependency Tree	3
1.3	Feed-forward neural network	4
2.1	CBOW	10
2.2	Skip-Gram	10
2.3	Weighting function f with $\alpha = 3/4$	11
2.4	GloVe vs Skip-gram	12
2.5	GloVe vs CBOW	12
3.1	Recurrent Neural Network	14
3.2	Long-Short Term Memory	16
3.3	Gated Recurrent Unit	17
3.4	Recursive Neural Network	20
3.5	Simple Recursive Neural Network	21
3.6	Syntactically Untied Recursive Neural Network	21
3.7	Matrix-Vector Recursive Neural Networks	22
3.8	Negated positives	22
3.9	Negated Negative	23
3.10	X but Y conjunction	23
3.11	Recursive Neural Tensor Network	24
3.12	Tree-Based LSTM Memory Cell Composition	25
3.13	Tree-Based GRU Hidden State Composition	27
4.1	Vanilla vs Momentum	32
4.2	Momentum vs Nesterov Momentum	32
4.3	Standard Feed-forward Neural Network	36
4.4	After applying Dropout	36
5.1	Binary Classification Average Training Time	38
5.2	Binary Classification Average Training Loss	38
5.3	Binary Classification Training Process	39
5.4	Fine Grained Classification Average Training Time	39

5.5	Fine Grained Classification Average Loss	39
5.6	Fine Grained Classification Training Process	40

LIST OF TABLES

3.1	Negations	24
5.1	Sentiment Classification Accuracy	40
5.2	Sentiment Classification Standard Deviation	41

INTRODUCTION

1.1 Motivation

This paper aims to contribute to the Artificial Intelligence community by proposing a different architecture (Tree-based GRU 3.2.6) and presenting its results in comparison to Tree-Based LSTM [77] (3.2.5). The conduction of this experiment is done as GRU architecture is less complicated and has less parameters to compute than LSTM. Therefore, we hypothesize that it is faster to train. GRU is a new approach, it is not determined whether it is better than LSTM or not [10], so far it is assumed that the comparison between those two models is like a comparison between non-linear activation functions, there is no "best" function but some non-linearities suit better some problems.

1.2 Overview of Thesis

This thesis consists of 5 chapters. The first chapter, the introduction, provides the reader with the necessary background to be able to read this report. The second chapter (2), word embeddings, is about the representation of the words on the vector space. On the second chapter we will go through the word embeddings benefits and the most efficient algorithms that are used to achieve this transformation (from words, to vectors). It is important for the reader to be aware of the fact that word embeddings are the input for the Tree-based GRU (and every other language model that we cover in this report). The third chapter (3), feedback neural networks, makes an in depth commentary of the feedback neural networks and its most common architectures. The forth chapter (4), network training, goes through the gradient descent algorithm, its different variants and extensions but also covers the properties of the neural network's

hyperparameters. Finally, the fifth chapter(5), experiments, demonstrates the comparison between Tree-based GRU and Tree-based LSTM, it also describes the gradient descent variants and the hyperparameters that were selected for the training of the model.

1.3 Background

This section provides the essential background on sentiment classification, sequences, syntactic structures, and neural networks.

1.3.1 Sentiment Classification and Sequences

Sentiment analysis/classification [59] (also known as opinion mining) is the classification on whether a piece of text is positive, negative or neutral using NLP, statistics, or machine learning methods.

A sequence is a string of objects. Each sequence is a row of item sets. The individual elements in a sequence are also called terms. In the case of the word sequences, a word corresponds to an item. In the case of the health care data, a test value is an item. Treating data as sequential (when they are sequences) improve the prediction accuracy of the classifiers [50].

1.3.2 Syntactic Structure

The sentences can be described with two ways for a machine to make sense out of them, the first one is by breaking up the sentence to phrases(constituents), which is known as constituent structure, and the second is by connecting the words with links, which is known as dependency structure. Those structures are constructed by parsing algorithms. The parsers for constructing a constituent tree are called phrase structured parsers and the parsers for dependency structures are called dependency parsers.

1.3.2.1 Constituency Structure

The phrase structure was introduced by Noam Chomsky, the idea behind constituency phrase structure is to organize the words into nested constituents[58] [46](Figure 1.1). Meaning that each nested constituent (word phrase) is a word unit. There has been different criteria for determining the constituents. The most popular is the one that claims that a constituent behaves as a unit no matter the place that is located in the sentence. Regarding to the construction of this structure, one great parser is the constituent parser from Zhu[83].

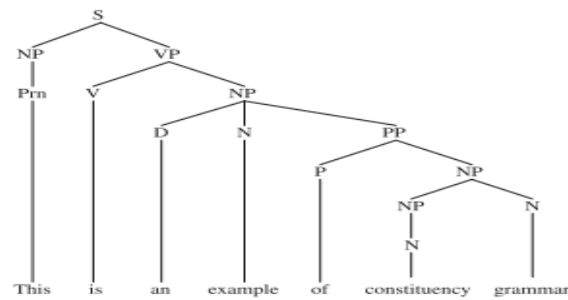


Figure 1.1: Constituency Tree

Source :“A Generative Constituent-Context Model for Improved Grammar Induction.”, D.Klein, 2002

1.3.2.2 Dependency Structure

The idea behind dependency structure is having words connected with a dependency relation (Figure 1.2), where one of them is the head and the other is the dependent and there is a link connecting them. In more detail, the dependent is the modifier, object, or complement while the head determines the behavior of the pair. The dependent requires the presence of the head; the head on the other hand doesn't require the presence of the dependent. [17]. In general, the dependency structure is a tree with the main verb as its root (head of the whole structure). It is worth mentioning that a dependency structure can be constructed from constituent trees as well [19]. The idea behind dependency structure is to directly show for the words of a sentence which are the words depend on (modify or are arguments of) which other words. [45]

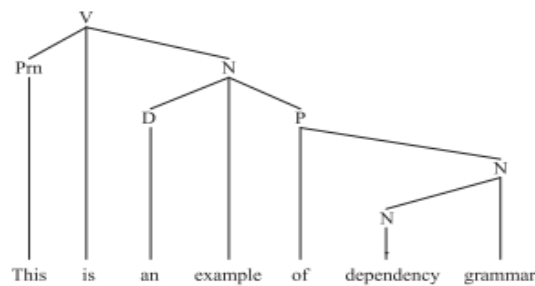


Figure 1.2: Dependency Tree

Source :“A Generative Constituent-Context Model for Improved Grammar Induction.”, D.Klein, 2002

1.3.3 Neural Networks

Neural networks are models of computation that were inspired by the way (we assume) our brain works [52], [65], [66]. The structure of artificial neural networks (ANN) is a network of small processing units (neurons) that are connected with weighted joints. Over the years many variants of ANN has been proposed. One important distinction is the way they are connected, with cycles or without. The former case of neural networks

are called feedback or recursive neural networks and will be examined at chapter 3, the latter case of neural networks are the Feedforward networks that will be examined at the next subsection.

1.3.3.1 Feedforward Networks

Given the absence of cycles, all nodes can be arranged into layers, and the outputs in each layer can be calculated given the outputs from the lower layers. The input i to a feedforward network is provided by setting the values of the lowest layer. Each higher layer is then successively computed until output is generated at the output layer o (Figure 1.3).

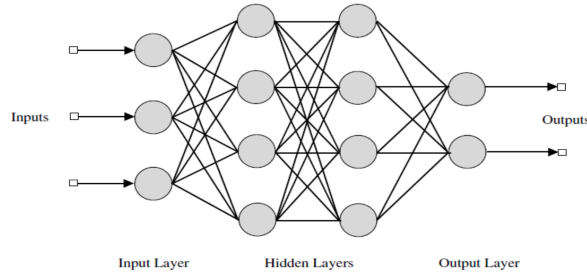


Figure 1.3: Feed-forward neural network

Source :

<https://www.linkedin.com/pulse/learning-scale-end-then-logic-ishtiaq-rahman>

Let w_{jk}^l be the weight for the connection from the k^{th} neuron in the layer $l-1$ to the j^{th} neuron in the l^{th} layer, b_j^l the bias of node j at layer l and α_j^l for the activation of neuron j at l^{th} layer. The equation below is using the sigmoid function.

$$\alpha_j^l = \sigma \left(\sum_{limit=k} w_{jk}^l \alpha_k^{l-1} + b_j^l \right) \quad (1.1)$$

Having the equation above in mind, we can rewrite it in a more compact and vectorized form:

$$\alpha^l = \sigma(w^l \alpha^{l-1} + b^l) \quad (1.2)$$

Let $z^l = \sum_{limit=k} w_{jk}^l \alpha_k^{l-1} + b_j^l$ be the weighted input to the neurons in layer l .

$$\alpha^l = \sigma(z^l) \quad (1.3)$$

The most popular choices of activation function are the hyperbolic tangent (1.4), sigmoid (1.5), rectified linear unit(1.6), softmax (generalization of sigmoid for K classes)(1.7).

$$\tanh(x) = \frac{e^{2x} - 1}{e^{2x} + 1} \quad (1.4)$$

$$\sigma(x) = \frac{1}{1 + e^{-x}} \quad (1.5)$$

$$f(x) = \max(0, x) \quad (1.6)$$

$$\text{softmax}(x)_j = \frac{e^{x_j}}{\sum_{k=1}^K e^{x_k}}, \text{ for } j = 1, \dots, K \quad (1.7)$$

The most popular FNNs are perceptrons[65], Kohonen maps[63] and Hopfield nets[41] and multilayer perceptron (MLP)[66],[80], [6].

1.3.3.2 Backpropagation

The most successful algorithm for training neural networks is backpropagation, it was introduced for this purpose by Rumelhart, Hinton, Williams[66] and some alterations were suggested by Zipser[84], and Werbos [80]. Backpropagation uses the chain rule to calculate the derivative of the loss function L with respect to each parameter(weights and biases) in the network. The weights are then adjusted by gradient descent algorithm (which we will go into detail at chapter 4). While it is not certain that backpropagation will reach a global minimum (unless the loss surface is convex¹) ,many researchers have worked on heuristic pre-training and optimization techniques that make them practically good enough for supervised learning tasks.

To calculate the gradient in a feedforward neural network, backpropagation proceeds as follows. First,proceeds to the forward pass, an example is propagated forward through the network to produce a value α_j^l , at each node j at layer l and outputs α^L at the output layer L . Then, a loss function value $L(\alpha_k^L, y_k)$ is computed at each output node k . Subsequently, for each output node j , we calculate the error where the first expression $\frac{\partial L(\alpha_j^L, y_j)}{\partial \alpha_j^L}$ corresponds to the rate of change in respect to the output neuron j , and the second term measures how fast the activation function σ is changing at z_j^L :

$$\delta_j^L = \frac{\partial L(\alpha_j^L, y_j)}{\partial \alpha_j^L} \sigma'(z_j^L) \quad (1.8)$$

¹convex surface:when the local optima is equal to the global optima

The equation above could be written in a more compact and matrix-based form, where $\nabla_{\alpha} L$ corresponds to the rate of change of L with respect to the output activations, as:

$$\delta^L = \nabla_{\alpha} L \odot \sigma'(z^L) \quad (1.9)$$

Having computed the error of the output layer, we go to compute the error of the prior layer. The equation for computing the layer before is:

$$\delta^l = ((w^{l+1})^T \delta^{l+1}) \odot \sigma'(z^l) \quad (1.10)$$

By combining the equations 1.9 and 1.10, we can compute the error at any layer. The backpropagation algorithm starts from the output layer L calculating the δ^L with equation 1.9 and moves to the previous layers with equation 1.10.

The equation for the rate of change of the cost in respect to the bias of node j in layer l is:

$$\frac{\partial L}{\partial b_j^l} = \delta_j^l \quad (1.11)$$

A more clean and vectorized form of the equation above can be written as:

$$\frac{\partial L}{\partial b} = \delta \quad (1.12)$$

The equation for the rate of change of the cost in respect to the weight that connects the node k and node j in layer l is:

$$\frac{\partial L}{\partial w_{jk}^l} = \alpha_k^{l-1} \delta_j^l \quad (1.13)$$

where α is the activation of the neuron input to the weight w and δ is the error of the neuron output from the weight w .

$$\frac{\partial L}{\partial w} = \alpha^{IN} \delta^{OUT} \quad (1.14)$$

CHAPTER 2

WORD EMBEDDINGS

Word embedding is a representation of a word in vector space where semantically similar words are mapped to nearby points. Word embeddings can be trained and used to derive similarities between words. They are an arrangement of numbers representing the semantic and syntactic information of words in a format that computers can understand. For many years, NLP systems and techniques would represent meaning of words using WordNet (George A. Miller, Princeton University, 1985) which is basically a very large graph that defines different relationships between words. In vector space terms, every word is a vector with one 1 and a lot of zeros (vocabulary size -1). This is a so called one-hot that describes words in the simplest way. However, this discrete representation had many issues, such as missing nuances, missing new words, the requirement of human labor to create and adapt, it was hard to compute accurate word similarity and most importantly when the vocabulary is large, the vector representation is gigantic.

The new approach of representing words was inspired by the quote "You shall know a word by the company it keeps"[27]. Instead of representing a word by its own index, represent a word by means of its words. This approach of representing words is by creating a dense embedding vector (word embedding). The word embeddings are derived by Vector Space Models (VSM) [67] which are divided in two categories which have been critically compared by Baron [2]. The first type of models is the count-based method (aka full document method) that compute the statistics of how often some words co-occur with its neighbor words in a large text corpus and then map those statistics to a low dimensional, dense vector. Some worth mentioning examples of this methodology are Hyperspace Analogue to Language [9], COALS method [21], Hellinger PCA [13]. The most popular model of the count-based method is the Latent Semantic Analysis [20] which we will cover at the section (2.1). The second

type of VSM models are the predictive models which try to predict directly the word from its neighbors in terms of learned low dimensional, dense embedding vectors. Some models worth mentioning of this category are Semantic Role Labeling(SRL) [15], Mnih and Kavukcuoglu vLBL and ivLBL[55],[54],Levy [48] proposed explicit word embeddings based on a PPMI metric. At sections 2.2 and 2.3 we will cover the most popular models of the VSM predictive models , named Word2Vec and GloVe. The 2.1 and 2.2 are covered just to demonstrate way that GloVe 2.3 was conceived, therefore they are covered with not much details.

The new approach of word representation tackles the problems of high dimensionality, the scalability of the vocabulary and the semantic relatedness of the words.

2.1 Latent Semantic Analysis

Latent semantic analysis (LSA) is a methodology in natural language processing of analyzing relationships between a set of documents and the terms they contain by producing a set of concepts related to the documents and terms. LSA assumes that words that are close in meaning will occur in similar pieces of text. The term-document matrix that is created, is quite sparse. Therefore a mathematical technique called singular value decomposition (SVD) is used to reduce the number of rows while preserving the similarity structure among columns. Singular Value Decomposition [31] is a method for identifying and ordering the dimensions of the observation that exhibit the most variation. Once we have found where the most variation lies, we can find the best approximation of the original observation using fewer dimensions. Therefore, it can be used for dimensionality reduction tasks.

Let X be a matrix with m terms and n documents where element (i, j) describes the occurrence of term i in document j co-occurrence matrix. According to SVD, there is a decomposition of X so that U and V are orthogonal matrices and Σ is a diagonal matrix. The values $s_1, \dots, s_l$ are called the singular values, and $u_1, \dots, u_l$ and $v_1, \dots, v_l$ the left and right singular vectors.

$$\begin{matrix} & X & & U & & \Sigma & & V^T \\ \begin{bmatrix} x_{11} & \cdots & x_{1n} \\ \vdots & \ddots & \vdots \\ x_{m1} & \cdots & x_{mn} \end{bmatrix} & = & \begin{bmatrix} u_{11} & \cdots & u_{1r} \\ \vdots & \ddots & \vdots \\ u_{m1} & \cdots & u_{mr} \end{bmatrix} & \begin{bmatrix} s_{11} & \cdots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \cdots & s_{rr} \end{bmatrix} & \begin{bmatrix} u_{11} & \cdots & u_{1r} \\ \vdots & \ddots & \vdots \\ u_{n1} & \cdots & u_{rn} \end{bmatrix} \end{matrix}$$

Moreover the The matrix product XX^T gives us the the correlation between the terms over the set of documents and the $X^T X$ gives us the correlation between the documents over the set of terms.

$$\begin{aligned} XX^T &= (U\Sigma V^T)(U\Sigma V^T)^T = (U\Sigma V^T)(V^T \Sigma^T U^T) = U\Sigma V^T V \Sigma^T U^T = U\Sigma \Sigma^T U^T = U\Sigma^2 U^T \\ X^T X &= (U\Sigma V^T)^T (U\Sigma V^T) = (V^T \Sigma^T U^T)(U\Sigma V^T) = V \Sigma^T U^T U \Sigma V^T = V \Sigma^T \Sigma V^T = V \Sigma^2 V^T \end{aligned}$$

Since $\Sigma \Sigma^T$ and $\Sigma^T \Sigma$ are diagonal we can safely conclude that U are the eigenvectors of XX^T , V are the eigenvectors of $X^T X$ and both products have the same eigenvalues,

given by the entries of $\Sigma\Sigma^T$ or $\Sigma^T\Sigma$. From Frobenius norm¹[32] we can derive that by taking the k largest singular values (express the importance of every word), and their corresponding singular vectors, we get the rank k approximation to X with the lowest error. The word vectors of the corpus will be the k columns of the matrix.

$$\begin{matrix} X & \hat{U} & \hat{S} & \hat{V}^T \\ \begin{bmatrix} x_{11} & \cdots & x_{1n} \\ \vdots & \ddots & \vdots \\ x_{m1} & \cdots & x_{mn} \end{bmatrix} & \approx \begin{bmatrix} u_{11} & \cdots & u_{1k} \\ \vdots & \ddots & \vdots \\ u_{m1} & \cdots & u_{mk} \end{bmatrix} \begin{bmatrix} s_{11} & \cdots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \cdots & s_{kk} \end{bmatrix} \begin{bmatrix} u_{11} & \cdots & u_{1k} \\ \vdots & \ddots & \vdots \\ u_{n1} & \cdots & x_{kn} \end{bmatrix} \end{matrix}$$

While this method solves the problem of dimensionality, it underlies some problems as well. First of all, the computational cost increases quadratically as the size of the matrix increases. Moreover when new words appear SVD has to be run again from scratch. Finally, it is able to find similarities between words, but it cannot represent relationships.

2.2 Word2Vec

Word2Vec (Mikolov, 2013) is a predictive model that learns word embeddings on an online way. The main idea behind Word2Vec is to predict the surrounding words of every word in a window of length m , instead of capturing all the co-occurrence counts directly. It is simpler and faster than LSA and can easily add a new word to the vocabulary.

The objective function of predictive VSMs (aka Neural probabilistic language models) aim to maximize the average log probability of any context word given the current center word where t is the number of tokens, m the co-occurrence window and θ all the variables we optimize using stochastic gradient descent.

$$J(\theta) = \frac{1}{T} \sum_{i=1}^T \sum_{-m \leq j \leq m, j \neq 0} \log p(w_{t+j}|w_t) \quad (2.1)$$

For $p(w_{t+j}|w_t)$:

$$p(o|c) = \frac{\exp(u_o^T v_c)}{\sum_{w=1}^W \exp(u_w^T v_c)} \quad (2.2)$$

where o is the output word id, c is the center word id, u is the center word vector and v is the output vector. This objective function is not scalable and it takes much time to train when the vocabulary is large.

Those two models are great at constructing word vectors, but when the vocabulary becomes too large the updates at each iteration take too much time. This problem is

¹Frobenius norm: is matrix norm of an $m \times n$ matrix A defined as the square root of the sum of the absolute squares of its elements

solved by a method called Negative Sampling [53]. This idea suggests to take random samples from the vocabulary that do not appear on the context and minimizing their probability of occurring 2.3.

$$J(\theta) = \frac{1}{T} \sum_{t=1}^T J_t(\theta) \quad (2.3)$$

where:

$$J_t(\theta) = \log \sigma(u_o^T v_c) + \sum_{i=1}^k \mathbb{E}_{j \sim P(w)} [\log \sigma(-u_j^T v_c)] \quad (2.4)$$

This objective function maximizes the probability of $u_o^T v_c$ co-occur and minimizes the probability the words randomly selected co-occur.

The most popular Word2Vec variants are CBOW [54], which stands for continuous bag of words, and Skip-gram [53]. CBOW (Figure 2.1) predicts the current word based on the context. More precisely, it predicts the current word based on the n neighbours that occur before and n neighbours that occur after than this word. The (window size) inputs share the weights that connect them with the hidden layer, what happens to the hidden layer, is to take the mean of the input words and it passes along to the output. Skip-gram (Figure 2.2) is the opposite of CBOW, while CBOW uses the context to predict the middle words Skip-gram uses the middle word to predict the context (window size). For the most part, CBOW tends to be a useful technique for smaller datasets. On the other hand, skip-gram treats each context-target pair as a new observation, therefore it tends to perform better with larger datasets.

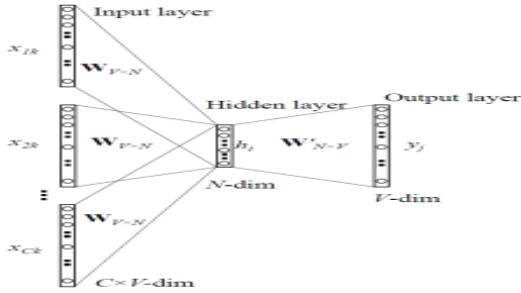


Figure 2.1: CBOW

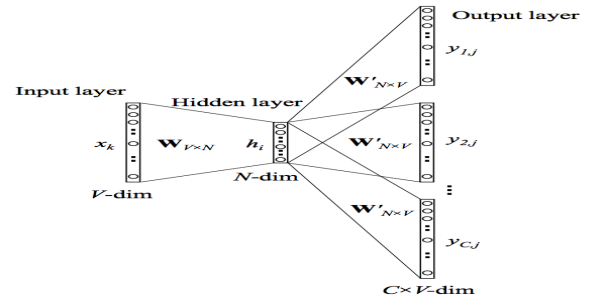


Figure 2.2: Skip-Gram

Source:

http://www.cs.nthu.edu.tw/shwu/courses/ml/labs/10_keras_Word2Vec/10_keras_Word2Vec.html

Finally according to Mikolov, Skip-gram works well with small amount of training data, and CBOW performs better in bigger amount of training data. All in all, Word2Vec models generate improved performance and are capable of capturing complex patterns beyond word similarity but they scale with corpus size and make inefficient usage of statistics.

2.3 GloVe

GloVe, which stands for global vectors, is an unsupervised learning algorithm for obtaining vector representations for words. Training is performed on aggregated global word-word co-occurrence statistics from a corpus, and the resulting representations are linear substructures of the word vector space [60].

GloVe combines the advantages of SVD and Word2Vec, it trains fast, it is scalable to huge corporas and performs well even with small corpus and small vectors. The main idea is, instead of going over one window at a time, it collects the whole corpus and is trained on the non-zero entries of the matrix. The weighted least squares regression model of GloVe is presented below 2.5, where W is the size of the vocabulary, $P_{i,j}$ is the probability that word j appear in the context of word i , $f()$ is the weighting function, u_i^T is the center vector of word i and v_j is the neighbor word j .

$$J(\theta) = \frac{1}{2} \sum_{i,j=1}^W f(P_{ij})(u_i^T v_j - \log P_{ij})^2 \quad (2.5)$$

The weighting function (Figure 2.3) should obey the following properties:

- $f(0) = 0$, If f is viewed as a continuous function, it should vanish as $x \rightarrow 0$ fast enough that the $\lim_{x \rightarrow 0} f(x) \log^2 x$ is finite.
- $f(x)$ should be non-decreasing so that rare co-occurrences are not overweighted.
- $f(x)$ should be relatively small for large values of x , so that frequent co-occurrences are not overweighted.

$$f(x) = \begin{cases} (x/x_{max})^a & \text{if } x < x_{max} \\ 1 & \text{otherwise} \end{cases}$$

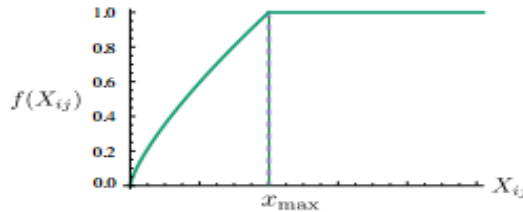


Figure 2.3: Weighting function f with $a = 3/4$.

Source:

<https://blog.acolyer.org/2016/04/22/glove-global-vectors-for-word-representation/>

The goal of the objective function 2.5 is to learn word vectors such that their inner product $(u_i^T v_j)$, which corresponding to the prediction of those words co-occurring, equals to the logarithm of the word's probability of co-occurrence $(\log P_{ij})$. The reason

behind using the logarithms of word's probability is the fact that the logarithm of a ratio equals the difference of logarithms, so it makes it feasible to illustrate probability of co-occurrence of two words as vector differences in the word vector space. For this reason, the resulting word vectors perform very well on word analogy tasks. (e.g. King - Man $\approx$ Queen - Woman)

Finally, the GloVe model compared to the two most popular Word2Vec models ,Skip-Gram (Figure 2.4) and CBOW (Figure 2.5),the plots illustrate the overall performance on the analogy task as a function of training time. Under the very same conditions GloVe outperforms word2vec. It trains faster and better irrespective of speed.

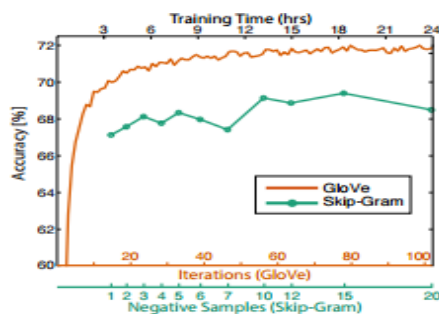


Figure 2.4: GloVe vs Skip-gram

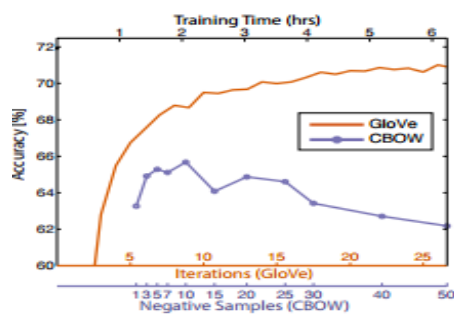


Figure 2.5: GloVe vs CBOW Source:

<https://blog.acolyer.org/2016/04/22/glove-global-vectors-for-word-representation/>

FEEDBACK NEURAL NETWORKS

This chapter aims to provide the reader with the background of feedback neural networks. In more detail, at section 3.1 we will review recurrent neural networks (simplistic subset of recursive neural networks) and at section 3.2 we will review recursive neural networks (also known as Tree-based recurrent neural networks).

3.1 Recurrent Neural Networks

A special case of recursive neural networks are the Recurrent Neural Networks (Figure 3.1) whose structure corresponds to a linear chain. The idea behind Recurrent Neural Networks is to make use of the sequential nature of the data. While in traditional neural networks we assume that input are independent to each other, in RNN we consider the importance of time. The reason that they are selected for sequential problems, is that they capture the information from each time-step, in practice though, they capture only a few steps but we will see the RNN variants that tackle this problem at subsections 3.1.2 and 3.1.3. Recurrent neural networks, has been under extensive research for quite some time now, so many variants of RNN have been proposed, such as Elman networks[25], Jordan networks[44], time delay neural networks[47], Long-Short Term Memory networks[39], echo state networks[42], Gated Recurrent Units[11]. In this section we will present Elman's networks(3.1.1, aka Simple recurrent networks), Long-Short Term networks(3.1.2), and Gated Recurrent Units(3.1.3).

3.1.1 Simple Recurrent Neural Networks

3.1.1.1 Forward Pass

The forward pass of an RNN is the same as that of an FNN, apart from the fact that the activations arrive at the hidden layer from the input layer (from the same time-step)

and from the hidden layer activations one time step before.

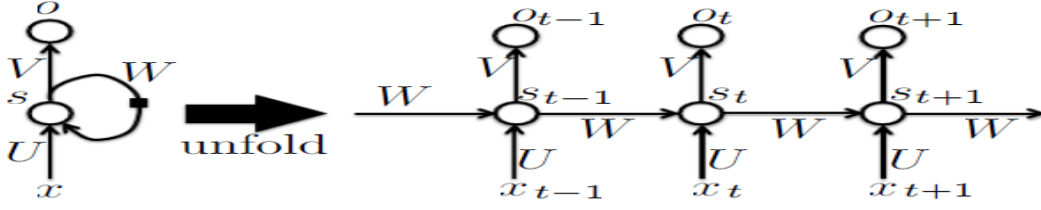


Figure 3.1: Recurrent Neural Network

Source: <http://www.wildml.com/2015/09/recurrent-neural-networks-tutorial-part-1-introduction-to-rnns/>

In more detail,(See Figure 3.1) x_t is the input at time step t , h_t is the hidden state at time step t 3.1(this is what differentiates FNN from RNN, h_t is the network's memory) ,and o_t is the output at time step t 3.2 .

(3.1)

(3.2)

Note that this requires initial values h_0 to be chosen , which correspond to the network's state before having any inputs. The initialization of the previous state can set the values to zero or initialize with a non zero initial values that are adjusted over the training process [24] which in some cases can improve the robustness and stability of the network.

3.1.1.2 Backward Pass

Regarding the RNN's training, the Backpropagation Through Time algorithm is applied [84][80]. BPTT is very similar to the Backpropagation (1.3.3) algorithm. The reason BPTT is used, is because it is more efficient in computing time. Moreover it is an algorithm under active research, one example is the new state of art way of computing BPTT has been published [34] which uses dynamic programming to balance a trade-off between caching of intermediate results and recomputation.

Just like Backpropagation (1.3.3) BPTT is consisted of repeated application of the chain rule. The only difference is that the objective function depends on the activation of the hidden layer not only because of its influence at output layer (o_t), but also on the hidden layer at the next time-step (h_{t+1} , which means that it affects the next steps to come as well). It is important to notice that the weights between time steps share the same values and that we sum over the whole sequence to get the derivatives with respect to each of the network weights. Mathematically:

$$\delta_{t,h} = \theta'(v_{t,h}) \left(\sum_{k=1}^K \delta_{t,k} u_{hk} + \sum_{l=1}^H \delta_{t+1,l} u_{hl} \right) \quad (3.3)$$

where,

$$\delta_{t,j} = \frac{\theta O}{\theta v_{t,j}} \quad (3.4)$$

The complete sequence of δ terms can be calculated by starting at last time step and recursively applying 3.3, decrementing t at each step. Finally to take the derivatives with respect to the network's weights, we sum over the whole sequence (since a single training instance is the full sequence).

$$\frac{\theta O}{\theta u_{ij}} = \sum_{t=1}^T \frac{\theta O}{\theta v_{tj}} \frac{\theta v_{tj}}{\theta u_{ij}} = \sum_{t=1}^T \delta_{t,j} b_{t,i} \quad (3.5)$$

where, $b_{t,i}$ is the activation of unit i at time t

While RNNs work well with sequences, they have a major drawback which is the vanishing/exploding gradient problem [38],[3]. One approach to deal with it is called the truncated BPTT [38] which is to set a limit on the time steps that you consider. Other approaches have been taken in order to tackle this problem. The most important were discrete error propagation [3], time delays [47], hierarchical sequence compression [68], Hessian Free Optimization [51], and Echo State Networks [43]. However, the most effective approach was a whole different architecture, the Long Short Term Memory (LSTM) architecture [39]. Moreover, RNNs have no control on "how many items/time steps to remember". The context needed for each classification task varies significantly, it may need 3 time steps (words) or 20. This problem is also solved with the LSTM architecture.

3.1.2 Long-Short Term Memory

Long-Short Term Memory (Figure 3.2) is the most popular architecture for sequential data, it is almost a subject under active research and constant improvement. LSTM networks have been applied to a variety of sequence modelling and prediction tasks with state-of-art performance, notably machine translation [1],[75], speech recognition [33], image caption generation [79], and program execution [81]. In its original form, LSTM contained only input and output gates. The forget gates [28], along with additional peephole weights [29] connecting the gates to the memory cell were added later. LSTMs are designed to be able to keep the important information and forget the noise. The default behavior of LSTM is to remember long periods and after training it remembers only the important information.

3.1.2.1 Architecture

Long-Short Term Memory networks' structure is like Simple RNN's, with only difference on the hidden layer structure. While RNNs have a single layer, the LSTMs have four hidden layers. The LSTM is consisted of cells and gates, cells contain the information and gates regulate "how much" of information to let go, gates are composed of a sigmoid layer that outputs an interval from 0 to 1 where 0 stands for "pass no information" and 1 for "pass everything". In more detail, it is consisted of the forget gate (f_t) which decides how much information (current and past) to pass to the network, the input gate (i_t) which decides what information to store at the cell, the cell state (C_t) which is the updated state, and finally the output gate.

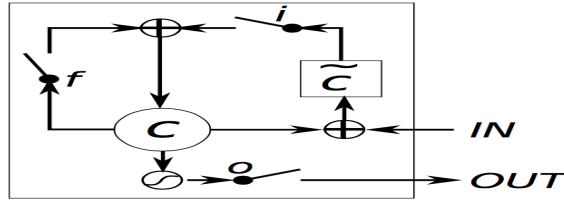


Figure 3.2: Long-Short Term Memory
Source : <https://deeplearning4j.org/lstm.html>

3.1.2.2 Forward Pass

The first step is to decide what to input to the network, here is where forget gate comes in (3.7) where takes as input the previous state (h_{t-1}) and the current input (x_t) and it outputs an interval between 0 to 1 to the candidate cell state ($\tilde{C}_t$) (3.8). Then we decide what values to update (3.9) at the input gate (i_t). After having computed candidate cell and the input gate we combine them with the previous cell state (C_{t-1}) to derive the current cell state (C_t) (3.10). Finally we feed the cell state to the output gate in order to decide how much of the cell state to output (3.11) as well as to a tanh layer so to squash the values between -1 and 1 (A.1).

$$i_t = \sigma(W^{(i)}x_t + U^{(i)}h_{t-1}) \quad (3.6)$$

$$f_t = \sigma(W^{(f)}x_t + U^{(f)}h_{t-1}) \quad (3.7)$$

$$o_t = \sigma(W^{(o)}x_t + U^{(o)}h_{t-1}) \quad (3.8)$$

$$\tilde{C}_t = \tanh(W^{(C)}x_t + U^{(C)}h_{t-1}) \quad (3.9)$$

$$C_t = f_t \odot C_{t-1} + i_t \odot \tilde{C}_t \quad (3.10)$$

$$h_t = o_t \odot \tanh(C_t) \quad (3.11)$$

3.1.2.3 Backward Pass

The original LSTM training algorithm [39] used an approximate error gradient calculated with a combination of Real Time Recurrent Learning [64] and Backpropagation Through Time (BPTT) [84]. The BPTT part was truncated after one time-step, since memory blocks would deal with the longer dependencies. The truncating method has the benefit of making the algorithm online, meaning that weight updates can be made after every time-step. However in this report we will extract the LSTM gradient with BPTT [Graves2005a], for further details look at appendix A.

3.1.3 Gated Recurrent Unit

Gated recurrent unit (Figure 3.3) is another variant of recurrent units proposed by KyungHyun Cho [10]. It is closely related Long-Short Term Memory. The GRU also controls the flow of information like the LSTM, but without using a memory unit. It exposes the full hidden content without any control.

3.1.3.1 Architecture

A GRU has two gates, a reset gate r , and an update gate z . The reset gate indicates how to combine the new input with the previous memory. The update gate defines how much of the previous state to keep. The basic idea of using a gating mechanism to learn long-term dependencies is the same as in a LSTM, but there are a few differences in terms of architecture. First of all, it doesn't have an output gate so it has fewer parameters (two gates, instead of three). Secondly the input and forget gates are substituted by an update gate z and the reset gate r is applied directly to the previous hidden state. Thus, the responsibility of the reset gate in a LSTM is really split up into both r and z . Finally we don't apply a second nonlinearity when we compute the output.

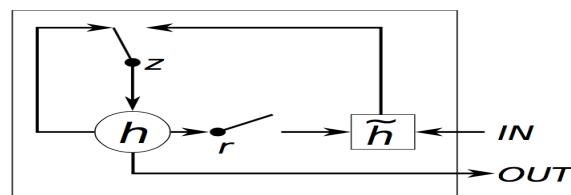


Figure 3.3: Gated Recurrent Unit
Source : <https://deeplearning4j.org/lstm.html>

3.1.3.2 Forward Pass

The first step is to determine the input values, this is update gate's task it will decide how much of the previous hidden state and how much of the candidate hidden state

combines to get the new hidden state (3.12). After the update gate, the reset gate while it has the exact same functional form (3.13) as the update gate and all the weights are at the same size, what makes it different is its position at the model. The reset gate is multiplied by the previous hidden state it controls how much from the previous hidden state we will consider when we create the new candidate hidden state. In other words, it has the ability to reset the hidden state, if we set the reset gate to 0 from (3.14) we start over from a new sequence as if h_t is the beginning of a new sequence. However this is not the full picture since $\tilde{h}_t$ is just a candidate for the hidden state, the actual hidden state will be a combination of previous hidden state h_{t-1} and the candidate hidden state $\tilde{h}_t$ controlled by the update gate z_t (3.15).

$$z_t = \sigma(W^{(z)}x_t + U^{(z)}h_{t-1}) \quad (3.12)$$

$$r_t = \sigma(W^{(r)}x_t + U^{(r)}h_{t-1}) \quad (3.13)$$

$$\tilde{h}_t = \tanh(W^{(h)}x_t + U^{(h)}(h_{t-1} \odot r)) \quad (3.14)$$

$$h_t = (1 - z) \odot \tilde{h}_t + z \odot h_{t-1} \quad (3.15)$$

3.1.3.3 Backward Pass

Gated Recurrent Units also use the BPTT algorithm in order to be trained. We will derive the gradients for E (error), W, U and by hand using the chain rule, for further details look at appendix B

3.2 Recursive Neural Networks

Recursive Neural Networks (RNNs) [71], [61], [30], [16], [36] are perfect for settings that have nested hierarchy and an intrinsic recursive structure [73]. The syntactic rules of language are highly recursive, therefore we use that recursive structure with a model that complies with that property. It is important to notice that RNN don't comprehend sentences as sequences but as hierarchies which makes them ideal for semantic representation tasks (paraphrase detection [71], relation classification, sentiment analysis, phrase similarity) but they can't predict future items from a sequence (next word from a given sentence), something that Recurrent Neural Networks are very good at due to their linear structure. Recurrent Neural Networks represent sentences as parse trees (Figure 3.4).

What RNNs are very good at, is handling negation. Due to their hierarchical structure, when negation is spotted, the meaning is just being reversed. You have a label at every node of the tree, and the leaves of the tree represent words.

Moreover another reason that RNN are so popular for natural language processing tasks, is that the input sequence length (sentence in our case) is not a restriction, it can take inputs of arbitrary lengths. The latter benefit is accomplished by making the input vector of sentence a predefined size no matter the length of the sentence. ([4],[35], [14]. Essentially what RNNs do is to merge the semantic understanding¹ of the words, then the grammatical understanding² of the phrase or sentence, which results to a parse tree representation of a phrase or a sentence. Having understood the words, and knowing the way words are put together we can retrieve the meaning of the sentence. Even though grammatical understanding is an assumption and it is not proven that it improves the accuracy, it is still under debate but we will assume that it helps the model.

In short, it extracts from the sentence the syntactic structure, which indicates the relationship between phrases, and it identifies the meaningful phrases within the sentences and the relationship between them. In order to extract the vector representation of the sentence, the idea is to recursively merge pairs of representations of smaller segments to get representation that covers bigger sentences.

The Recursive Neural Networks are trained with the Backpropagation Through Structure algorithm[30] which is very similar to the standard Backpropagation, we use the 1.10 and the 1.14, that was discussed at subsection 1.3.3 with three minor differences. Firstly we sum up the derivatives of W from all the nodes, secondly we split the derivatives at each node and finally we add different error messages from parent node and self node.

Finally, it is assumed that the tree structured is given, which indicates that some preprocessing is required if it is not given. In our experiment we will use the Stanford Sentiment Treebank(SST) that was trained with the Stanford Parser[70], which is similar to max-margin parsing [78], to derive the tree structure is for every sentence. We will not go through the way that the sentence trees were constructed because we will add unnecessary complexity that is beyond the scope of this report, but at a high level explanation the parser that is used, have a loss term that penalizes the not plausible phrases.

3.2.1 Simple Recursive Neural Network

This model (Figure 3.5) is the standard recursive neural network. The first step to take is to take a sentence parse tree and the sentence word vectors and begin from the bottom leaves to the top root of the tree. The mathematical formula to merge

¹semantic understanding: understanding of the meaning of a sentence, represent accurately the phrase as a vector in a structured semantic space

²grammatical understanding: it is identified the underlying grammatical structure of the sentence

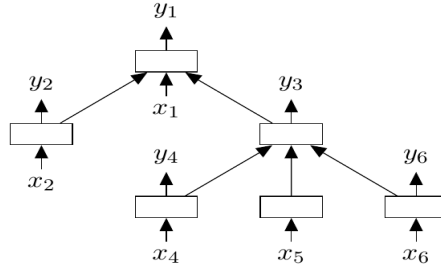


Figure 3.4: Recursive Neural Network

Source : <https://stats.stackexchange.com/questions/153599/recurrent-vs-recursive-neural-networks-which-is-better-for-nlp>

those two vectors (aka children) and create a new "word phrase"vector (parent) can be illustrated below (3.16). The h vector now represent the "this assignment"phrase. Having computed the vector representations of the sentences, we compute a s score 3.17 which represents the quality of the merge and decides which pair of representations to merge first. In order to derive some meaning of the word vector, we feed it to a softmax layer (3.18) to compute the score over a set of sentiment classes, a discrete set of known classes that represent some meaning. This process happens till the model reach the root of the tree. Moreover it is important to note that the W parameters is the same for all the nodes of the trees. It is obvious that this is quite a naive approach and linguistic complexity is higher than that. It is too much to ask from a simple function like this to capture the language complexity.

$$h_1 = \tanh\left(W \begin{bmatrix} c_1 \\ c_2 \end{bmatrix} + b\right) \quad (3.16)$$

$$s_1 = W^{score} p \quad (3.17)$$

$$\hat{y} = \text{softmax}(Wh_1 + b) \quad (3.18)$$

3.2.2 Syntactically Untied SU-RNN

One extension to the Simple RNN, the Syntactically Untied RNN model [73](Figure 3.6) was introduced to solve the problem mentioned at the previous subsection. What this model does, is to have unique weight matrices for every syntactic category. The syntactic categories are identified from the parser that determined the structure of the

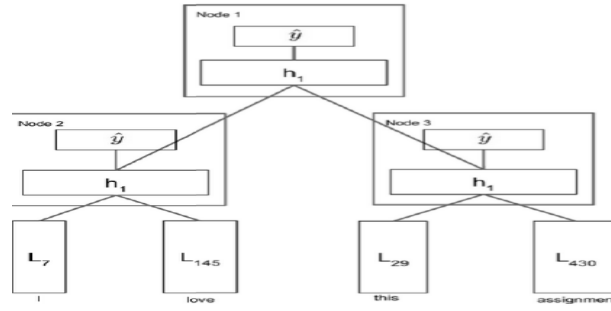


Figure 3.5: Simple Recursive Neural Network

Source: <https://www.slideshare.net/jiessiecao/parsing-natural-scenes-and-natural-language-with-recursive-neural-networks>

tree. This has proven to increase the weight matrices to learn and outperforms the methods that were mentioned till that point, but the performance boost we gained is not significant.

One impressive accomplishment of this model is that the trained weight matrices are capable of learning the semantics of the phrases. For example a determiner followed by a noun phrase (e.g. "an elephant") emphasizes more on the noun phrase than on the determiner. The architecture of the SU-RNN model compared to the Simple RNN is illustrated at Figure 3.6.

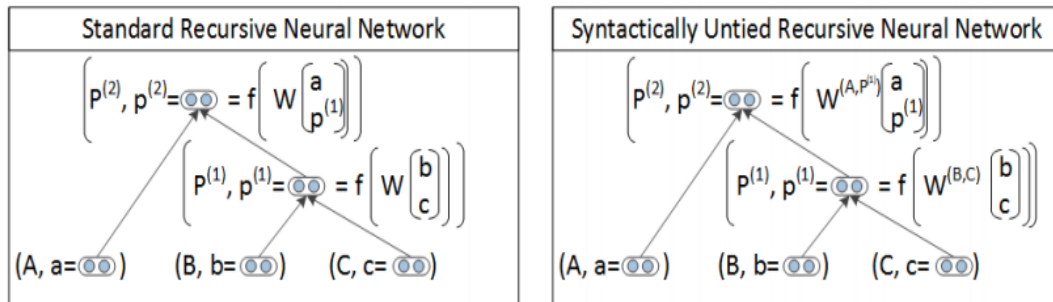


Figure 3.6: Syntactically Untied Recursive Neural Network

Source:

<https://wugh.github.io/posts/2016/05/cs224d-notes5-recursive-neural-networks/>

3.2.3 Matrix-Vector Recursive Neural Networks

Another alteration of Recursive Neural networks is the Matrix-Vector Recursive Neural Networks [72] (Figure ??) which improves the semantic representation of the sentences. The major difference is that not only we include a word vector (d -dimensional), but also a word matrix ($d \times d$) (Figure 3.7).

$$h_1 = \tanh\left(W \begin{bmatrix} C_2 c_1 \\ C_1 c_2 \end{bmatrix} + b\right) \quad (3.19)$$

This approach not only represents the meaning of each word but also the effect that it has on the neighboring words. Suppose we feed two words to the model, a and b, the parent vector is the concatenation of the word vector of the former multiplied with the word matrix of the latter and the word vector of the later is multiplied by the word matrix of the former (Ab and Ba). In the figure's example, the word matrix of "very" could also be the identity³ multiplied by a scalar (above one) which indicates the impact it has to the word "bad".

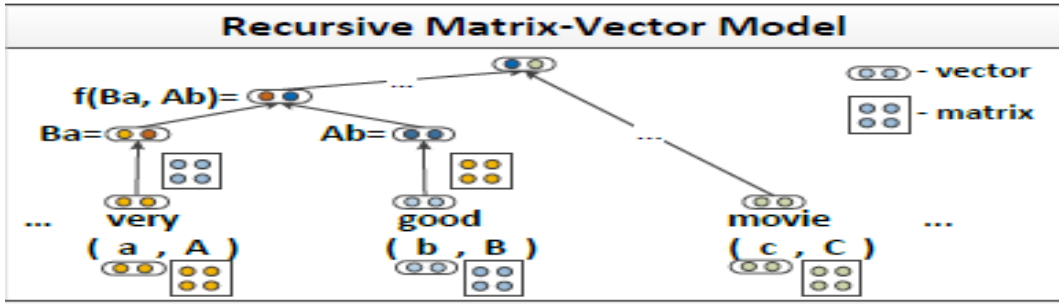


Figure 3.7: Matrix-Vector Recursive Neural Networks

Source:

<https://wugh.github.io/posts/2016/05/cs224d-notes5-recusive-neural-networks/>

Despite the fact that this is the most expressive model we have explored till now, it is still not good enough. It fails to capture the semantics of some relations. There have been observed three types of errors. [73] The first type (Figure 3.8) is the negated positives, this case occurs when something is classified as positive but one word turns it negative, the model can not capture the importance of that one word strong enough to flip the sentiment of the entire sentence.

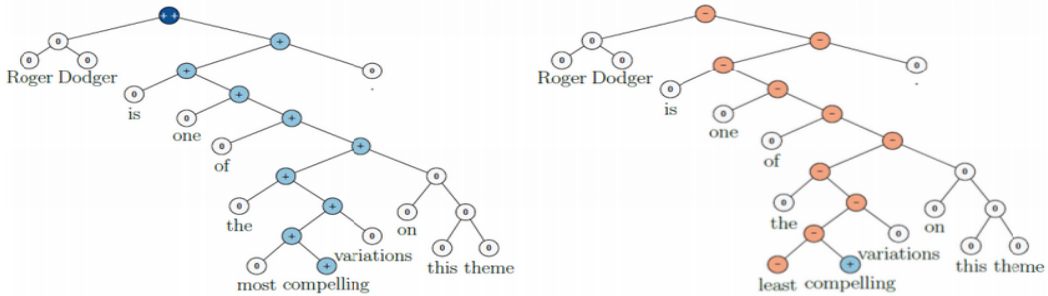


Figure 3.8: Negated positives

Source: https://cs224d.stanford.edu/lecture_notes/LectureNotes5.pdf

The second type (Figure 3.9) is the negated negative, where we say something is not

³identity matrix: square matrix with ones on the main diagonal and zeros elsewhere

bad. The MVRNN can not recognize that the word "not" because it turns sentiment from negative to neutral.

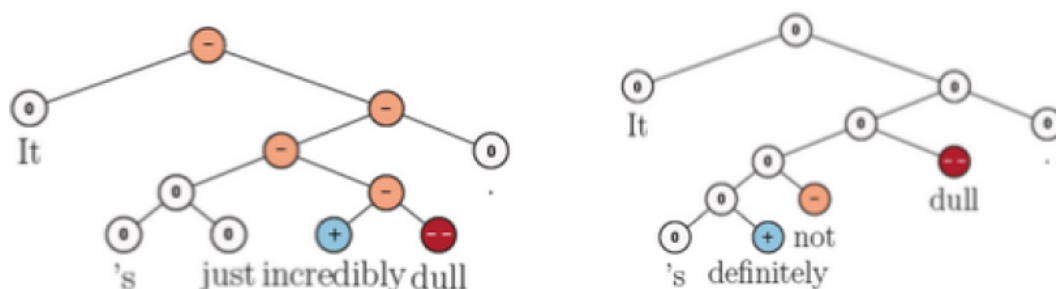


Figure 3.9: Negated Negative

Source: https://cs224d.stanford.edu/lecture_notes/LectureNotes5.pdf

The final type of errors (Figure 3.10) we observe is the "X but Y conjunction". In our example, X is negative but the Y is positive and the sentiment is positive. The MV-RNNs have some issues with such cases.

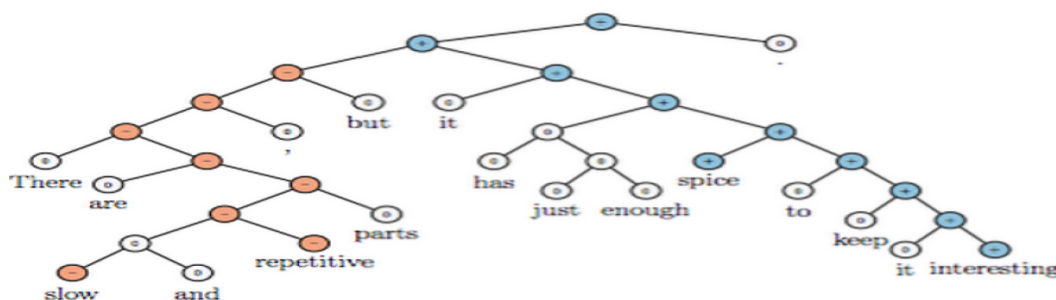


Figure 3.10: X but Y conjunction

Source: https://cs224d.stanford.edu/lecture_notes/LectureNotes5.pdf

Thus, we must look for an even more expressive composition algorithm that will be able to fully capture these types of high level compositions.

3.2.4 Recursive Neural Tensor Network

The third Recursive Neural Network variant that will be covered is the Recursive Neural Tensor Network (Figure 3.11). RNTN was conceived by Richard Socher [73] in order to solve the three types of errors we left of with at the subsection 3.2.3. Moreover it is famously quite successful for dealing with double negations. The Recursive Neural Tensor Network gets rid of the concept of a word matrix as well as the affine transformation⁴ pre-tanh/ σ concept that we saw before. To combine two word vectors or phrase vectors, we again concatenate them to form a vector $\in 2d$ but instead of putting it through an affine function then a nonlinear, we put it through a quadratic first, then a nonlinear, such as:

⁴affine transformation: is a transformation composed of a linear function+ a constant

$$h(1) = \tanh(x^T V x + W x) \quad (3.20)$$

Where V is a 3rd order tensor $\in^{2d \times 2d \times d}$. The quadratic shows that we can indeed allow for the multiplicative type of interaction between the word vectors without needing to maintain and learn word matrices. Figure 3.11: One slice of a RNTN. Note there would be d of these slices.

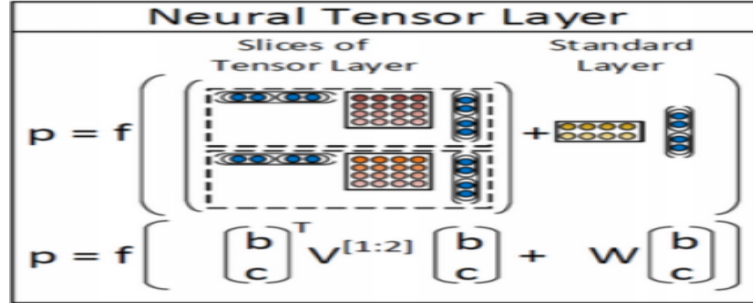


Figure 3.11: Recursive Neural Tensor Network

Source:

<https://wugh.github.io/posts/2016/05/cs224d-notes5-recusive-neural-networks/>

A major problem of the models we covered before is their inability to handle negation[69]. The table 3.1 below shows how the RNTN handles negations.

Table 3.1: Negations

Model	Negated Positive	Negated Negative
RNN	33.3	45.5
MV-RNN	52.4	54.6
RNTN	71.4	81.8

3.2.5 Tree-Based Long-Short Term Memory Networks

Tree-Based LSTM (Figure 3.12) was recently conceived by Kai Sheng Tai [77]. This is a hybrid model that combines LSTMs and Recursive neural network. It is important to notice that LSTMs were used for linear chained structured recursive neural networks (recurrent neural networks). The main difference that this model has with the standard LSTMs is that it is required the average of the child vectors and a special forget gate for each child. The idea behind this architecture is mostly to handle negation, by keeping in memory the semantically important words and forgetting the non significant. This process happens as the model goes through the tree structure. The Figure 3.12 illustrates how the new memory cell c_1 and hidden state h_1 are composed with two children.

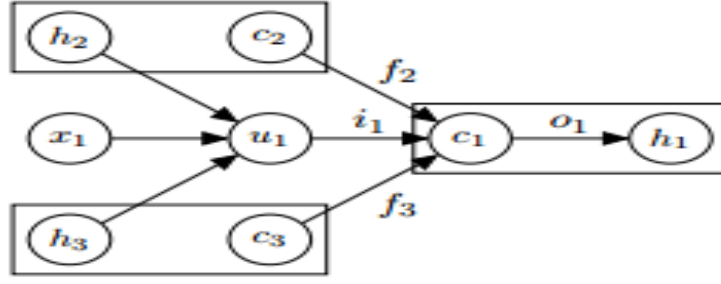


Figure 3.12: Tree-Based LSTM Memory Cell Composition
Source: <https://arxiv.org/pdf/1503.00075.pdf>

The Tree-LSTM behaves very similar to the standard LSTM, it takes as input vector x_j with only difference that the input vector depends on the tree structure (1.3.2). If the tree is a constituency tree, the leaf nodes take the corresponding word vectors as input, if the tree is a dependency tree each node in the tree takes the vector corresponding to the head word as input. The two extensions that Tai(2015) proposed are the Child Tree-LSTM and N-ary Tree-LSTM.

3.2.5.1 Child-Sum Tree-LSTM

Sum Tree-LSTM unit conditions its components on the sum of child hidden states h_k , this model performs well with high branching factor tree structures or with structures that its children are not ordered. Dependency trees is a good choice of structure for that model since the number of dependents of a head can be quite variant. Let $C(j)$ the set of children of node j and $k \in C(j)$, the equations of the model are the following:

$$\tilde{h}_j = \sum_k h_k \quad (3.21)$$

$$i_j = \sigma(W^i x_j + U^i \tilde{h}_j + b^i) \quad (3.22)$$

$$f_{jk} = \sigma(W^f x_j + U^f h_k + b^f) \quad (3.23)$$

$$o_j = \sigma(W^o x_j + U^o \tilde{h}_j + b^o) \quad (3.24)$$

$$\tilde{C}_j = \tanh(W^u x_j + U^u \tilde{h}_j + b^u) \quad (3.25)$$

$$C_j = i_j \odot \tilde{C}_j + \sum_{k \in C(j)} f_{jk} \odot C_k \quad (3.26)$$

$$h_j = o_j \odot \tanh(C_j) \quad (3.27)$$

The matrix parameters on the equations above can be interpreted as encoding correlations between the input x_j , the hidden states h_k of the children and the component vectors of the Tree-LSTM. A natural extension of this model is the Dependency Tree-LSTM which is the application of the Child-Sum Tree-LSTM to a dependency tree. This model has proven not to perform that well, I assume that this can be a result of its simplistic architecture and its lack of use of the sentiment contained at the leaves.

3.2.5.2 N-ary Tree-LSTM

On the other hand, the N-ary Tree-LSTM perform well where the branching factor is less or equal to N and the children are ordered (constituent). For any j node, the hidden state is h_{jk} and its memory cell is c_{jk} . The equations of the model are the following :

$$i_j = \sigma(W^i x_j + \sum_{l=1}^N U_l^i h_{jl} + b^i) \quad (3.28)$$

$$f_{jk} = \sigma(W^f x_j + \sum_{l=1}^N U_{kl}^f h_{jl} + b^f) \quad (3.29)$$

$$o_j = \sigma(W^o x_j + \sum_{l=1}^N U_l^o h_{jl} + b^o) \quad (3.30)$$

$$\tilde{C}_j = \tanh(W^u x_j + \sum_{l=1}^N U_l^u h_{jl} + b^u) \quad (3.31)$$

$$C_j = i_j \odot \tilde{C}_j + \sum_{l=1}^N f_{jl} \odot c_{jl} \quad (3.32)$$

$$h_j = o_j \odot \tanh(C_j) \quad (3.33)$$

It is important to notice that the N-ary Tree-LSTM model have separate parameter matrices for each child, which results in better learning of fine-grained conditioning on the states of a unit's children than the ChildSum Tree-LSTM. It can be considered as a case of trade-off of performance and computational cost. Suppose that an example of a constituent tree that its left child of a node is a noun phrase, and the right child a verb phrase. It is beneficial for this case to magnify the verb phrase. N-ary tree is able to do that by training the U_{kl}^f parameters so that the components of f_{j1} are close to 0 ("forget") while the components of f_{j2} are close to 1 ("remember").

A natural extension to N-ary LSTM model is the Constituency Tree-LSTM which is an application of Binary Tree-LSTM. This is the model that we will compared on the later section.

3.2.6 Tree-Based Gated Recurrent Unit

Tree-Based Gated Recurrent Unit (Figure 3.13) is a model that was inspired by the Tree-Based LSTM architecture. As far as I know it hasn't formally been published anywhere so I can't give the credentials to someone. The idea is almost the same with the Tree-based LSTM. However instead of having two forgetting gates, we will have two reset gates.

3.2.6.1 N-ary Tree-GRU

In this report, I have developed the alteration of N-ary Tree-LSTM named N-ary Tree-GRU. In our experiment, the tree-structure of the inputs is constituency trees the inputs will be word vectors. It is important for the reader to note that both architectures (Tree-LSTM and Tree-GRU) have an affect only on the composition of the parent node. Other than that, they are like the simple recursive neural network. In more detail, the N-ary Tree-GRU takes as input word vectors and it creates the candidate hidden state $\tilde{h}_j$ with the combination of the child nodes, its child has its own reset gate. After having the candidate hidden state $\tilde{h}_j$, we calculate the actual hidden state h_j .

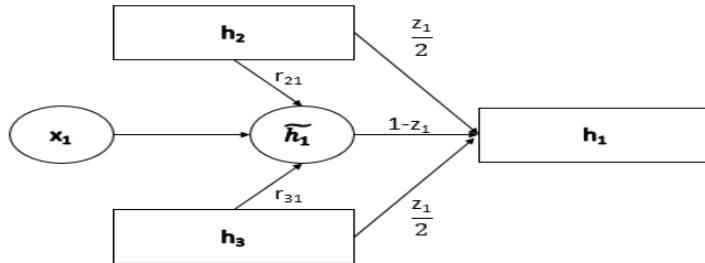


Figure 3.13: Tree-Based GRU Hidden State Composition

$$z_j = \sigma(W^z x_j + \sum_{l=1}^N U_l^z h_{jl} + b^z) \quad (3.34)$$

$$r_{jk} = \sigma(W^r x_j + \sum_{l=1}^N U_{kl}^r h_{jl} + b^r) \quad (3.35)$$

$$\tilde{h}_j = \tanh(W^h x_j + \sum_{l=1}^N U_l^h (h_{jl} \odot r_{jl}) + b^h) \quad (3.36)$$

$$h_j = (1 - z_j) \odot \tilde{h}_j + \sum_{l=1}^N \frac{z_j}{N} \odot h_{jk} \quad (3.37)$$

NEURAL NETWORK TRAINING

This chapter is devoted to the training process of a neural network which is consisted of the backpropagation and gradient descent algorithms[8]. We have discussed in great detail regarding the backpropagation algorithm at the subsection 1.3.3 and sections 3.1, 3.2 (BPTT and BPTS). However we haven't covered the gradient descent algorithm and its different extensions. Gradient descent algorithm is an algorithm under active research. In this chapter it will be covered the process of updating the network parameters and the impact of selecting the right hyperparameters. This chapter will be divided into three sections, the first section will cover the different variants of gradient descent algorithm data selection 4.1, the second section will go through the most common gradient descent extensions 4.2 and the third section will be about the neural network's hyperparameters 4.2 which are vital to the neural network's adequate training.

4.1 Gradient Descent Variants

On this section, there will be covered three different variants of the gradient descent algorithms with regards to the amount of data they take to compute the gradient of the objective function. On the subsection 4.1.1 will be covered the Batch Gradient Descent which is the maximization of the likelihood over the entire training set, which is quite slow for big data sets. The second alternative is the Stochastic Gradient Descent (4.1.2) which depends on the error of one particular sample, this variant makes the computation much faster but it has to proceed through many iterations in order to show indications of improvement. The mini-batch Gradient Descent (4.1.3) is the compromise between the two variants discussed before. [57]

4.1.1 Batch Gradient Descent

The Batch Gradient Descent calculates the gradient of the objective function with regards to the parameters θ of the entire training set(4.1). Since each update is performed after having calculated the whole training dataset, it is easy to conclude that it is a quite slow way of training in cases that data sets are large. Moreover another disadvantage of this method is that it doesn't allow to update the model on an online way.

$$\theta = \theta - \eta * \nabla_{\theta} J(\theta) \quad (4.1)$$

4.1.2 Stochastic Gradient Descent

On the other hand, Stochastic Gradient Descent performs a parameter update for each sample i (4.2)[7]. This variant much faster than the Batch Gradient Descent and it is able to learn in an online way. However, this approach has a few flaws as well. During the first steps of the training, the objective function fluctuates heavily because the updates have high variance. This characteristic has strengths and weaknesses as well, because it can land to a potentially better local minima, or it can make the minimization process too complicated. Despite its high variance at the beginning of the training, on the long run it will become more stable.

$$\theta = \theta - \eta * \nabla_{\theta} J_i(\theta) \quad (4.2)$$

4.1.3 Mini-Batch Gradient Descent

Mini-Batch Gradient Descent[49] is the happy medium between the two contradictory variants discussed above(4.1.1,4.1.2). This approach performs the parameter updates based on the gradient of the parameter from n samples (batch size) of the training set (4.3). It tackles the problem of update's high variance which leads to a more stable convergence and it trains faster than the vanilla gradient descent.

$$\theta = \theta - \eta * \nabla_{\theta} J_{i:i+n}(\theta) \quad (4.3)$$

Despite its advantages, Mini-Batch Gradient Descent underlies some flaws. A key challenge of minimizing highly non-convex error functions common for neural networks is avoiding getting trapped in their numerous suboptimal local minima. Dauphin[18] argues that the difficulty arises in fact not from local minima but from saddle points, i.e. points where one dimension slopes up and another slopes down. These saddle points are usually surrounded by a plateau of the same error, which makes it notoriously hard for SGD to escape, as the gradient is close to zero in all dimensions.

4.2 Gradient Descent Extensions

There are several approaches for performing the parameter update. In this section it will be covered the most popular gradient descent optimization algorithms that are used to tackle the challenges that were mentioned at section 4.1. Those techniques are used to solve some of the problems that the vanilla gradient descent algorithm underlies. Namely, the challenge of finding a right learning rate (4.3.1), the learning rate is the same for each parameter which can be a problem if the data is sparse, being trapped to a suboptimal local minima (4.1.3). The techniques that will be discussed are focused on the learning rate manipulation and its impact on the parameters' update.

4.2.1 Vanilla update

Vanilla update is the simplest form of update, it performs the updates of the parameters towards the negative gradient direction. So far, when gradient descent algorithm was mentioned we were referring to the vanilla gradient descent(4.1,4.2,4.3).

4.2.2 Momentum update

Momentum update [76] speeds up when the parameter has a consistent gradient and slows down when the gradient changes directions. This approach was inspired by the physical perspective of the problem. Just like a ball rolling down a hill, the steeper the slope the faster the ball rolls, based on the same idea the momentum update was conceived. The property of velocity is integrated by adding a fraction γ (momentum term) of the past update vector from the previous time step to the current update vector(4.4). After having calculated the update vector we integrate it into the parameter (4.5).

$$v_t = \gamma v_{t-1} + \eta \nabla_{\theta} J(\theta) \quad (4.4)$$

$$\theta = \theta - v_t \quad (4.5)$$

The outcome of using momentum is faster convergence and reduced oscillation[62]. At figure 4.1 we can see the difference between the vanilla gradient descent updates(on the left side) and the momentum updates(on the right side).

4.2.3 Nesterov Momentum update

Nesterov Momentum update is a smarter alteration of vanilla momentum update. The main idea of Nesterov Momentum is for the update to have a notion of where the gradient is heading therefore it slows down before the gradient changes direction. This technique is also known as Nesterov Accelerated Gradient (NAG) [56],[74] it makes

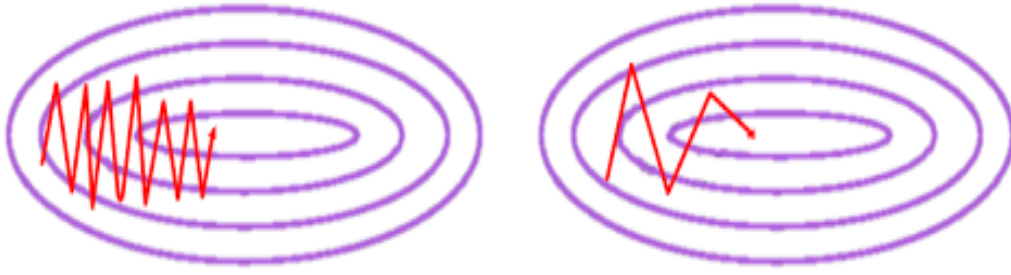


Figure 4.1: Vanilla vs Momentum

Source :

'<http://dsdeepdive.blogspot.com/2016/03/optimizations-of-gradient-descent.html>'

a rough approximation of where the parameter will be, so now it can effectively look ahead by calculating the gradient with regards to the approximation of the future position and not the current one (4.6). According to Bengio[5] the estimated update prevents us from going too fast and results in increased responsiveness.

$$v_t = \gamma v_{t-1} + \eta \nabla_{\theta} J(\theta - \gamma v_{t-1}) \quad (4.6)$$

$\theta = \theta - v_t$ Source : <http://dsdeepdive.blogspot.com/2016/03/optimizations-of-gradient-descent.html>

(4.7)

The difference between the two approaches can be illustrated below at Figure 4.2.

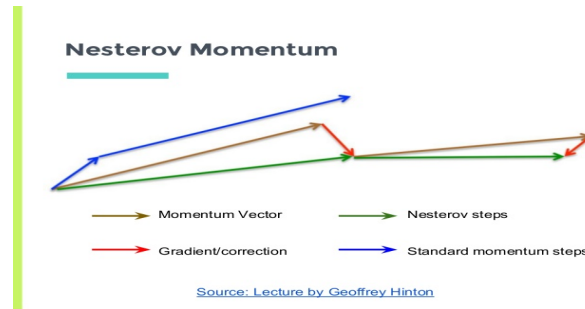


Figure 4.2: Momentum vs Nesterov Momentum

Source: <https://www.slideshare.net/cfregly/gradient-descent-back-propagation-and-auto-differentiation-advanced-spark-and-tensorflow-meetup-08042016>

4.2.4 AdaGrad

AdaGrad[23] adapts the learning rate for every feature which eliminates the need of manually tuning the learning rate. It adapts the learning rate to the parameters, performing larger updates for infrequent and smaller updates for frequent parameters. For this reason, it is well-suited for dealing with sparse data.[12] Moreover, Pennington [60] used Adagrad to train GloVe(section 2.3) word embeddings that we use at our

experiment on the next chapter. The reason he used AdaGrad, is because infrequent words require much larger updates than frequent ones.

Since AdaGrad uses different learning rate for every parameter θ_i at every time step t for the sake of convenience we assume $g_{t,i}$ (4.8) to be the gradient of the objective function of parameter i and g_t (4.9) be the vector of all the parameters at time step t . Apart from the element-wise property of this approach, Duchi makes use of a diagonal matrix $G \in R^{d \times d}$ that is consisted of the sum of the squares of the gradients with respect to the parameter up to time step t , while it also has a smoothing term ϵ that prevents it from being divided with 0. It is adapting the learning rate by caching the sum of squared gradients with respect to each parameter at each time step. The reason that we use the squared sums is not specified, but it performs way better than taking the matrix without the square root operation. Finally the formula for computing the parameter can be seen below.

$$g_{t,i} = \nabla_{\theta} J(\theta_i) \quad (4.8)$$

$$\theta_{t+1} = \theta_t - \frac{\eta}{\sqrt{G_t + \epsilon}} \odot g_t \quad (4.9)$$

Despite its high performance, Adagrad has a serious drawback which stems from the fact that the learning rate shrinks after some point due to its accumulation of the squared gradients in the denominator. This is the reason that it is notoriously aggressive at the machine learning community.

4.2.5 AdaDelta

In order to tackle the learning rate shrinking problem Zeiler came up with a different way of adapting the learning rate. Adadelta [82] comes to rescue. Adadelta is an extension of Adagrad that alleviates its aggressively monotonically decreasing learning rate. It was suggested that instead of accumulating all the previous gradients it would be better to set a fixed window of size w and take the accumulated gradients of size w . Moreover, another alteration of Adagrad is the way it treats the past gradients. Instead of storing the past squared gradients, the sum of gradients is defined as a decaying average of all past squared gradients. The average of time step t depends only on the previous average and the current gradient. It is also used the momentum term, that was covered at subsection 4.2.2, which determines how much off the past average will affect the current average.

$$E[g^2]_t = \gamma E[g^2]_{t-1} + (1 - \gamma) g_t^2 \quad (4.10)$$

The parameter update is defined similar to the AdaGrad but instead of the diagonal matrix G_t we use the decaying average of the past squared gradients $E[g^2]_t$.

$$\Delta\theta_t = -\frac{\eta}{\sqrt{E[g^2]_t + \epsilon}} g_t \quad (4.11)$$

Moreover the denominator of (4.11) is the root mean squared(RMS) error of criterion of the gradient, having noticed that Ziegel defined the decaying average of the squared parameter updates. However the $RMS[\Delta\theta]_t$ is not known so it is approximated with the RMS of the parameter updates from the previous time step. Therefore the updates of the parameters can be computed as :

$$\Delta\theta_t = \frac{RMS[\Delta\theta]_{t-1}}{RMS[g]_t} g_t \quad (4.12)$$

Something interesting about this approach is that the learning rate is irrelevant, as it is nowhere in the update rule.

4.2.6 RMSprop

RMSprop is an unpublished adaptive learning method from Hinton at his coursera [Machine Learning course](#) that is commonly used from the deeop learning community (it is even a [built-in function](#) at tensorflow). RMSprop is very similar with AdaDelta, in fact is is just like the first part of AdaDelta but instead of γ there is a default value of 0.9 while its suggested initial learning rate is 0.001. Even though they were established the same period they were conceived independently.

$$E[g^2]_t = 0.9E[g^2]_{t-1} + 0.1g_t^2 \quad (4.13)$$

$$\Delta\theta_t = -\frac{\eta}{\sqrt{E[g^2]_t + \epsilon}} g_t \quad (4.14)$$

4.3 Hyperparameters

4.3.1 Learning Rate

Learning rate can be thought as the rate that the parameter update based on its gradient. Choosing a proper learning rate can be difficult. A learning rate that is too small leads to slow convergence, while a learning rate that is too large can make the loss function fluctuate around the minimum or even diverge.

4.3.2 Regularization

It is a method for preventing overfitting. It essentially works by setting a penalty a complexity penalty to the loss function. In practice, this means that it penalizes a

function that is too non-linear and learns by heart the information that is contained in the training set and is not able to generalize well to new examples. In this paper we will mention the most popular regularizers that we will also use at the experiments. Let those be L2 regularizer and dropout regularizer. But why do we really need regularization? It doesn't help the model perform well at the training set but perform well at the new examples(test set), which means that it minimizes the generalization error. The generalization error is the sum of the squared bias and variance of the trained model. The variance of the model indicates how much the model varies if we change the training set, the bias of the model is how close is the model to the true solution (the model that generated those instances).

4.3.2.1 L2 Regularization

The L2 regularization method [40] adds a regularization term in order to prevent the coefficients to fit so perfectly to overfit. When it comes to neural networks it only regularizes the connection weights the hidden layers. In more detail, what we do is to penalize the square of the weight value for each hidden layer k and for each connection i, j . Notice that the sum of i, j from equation (4.15) corresponds to the Frobenius Norm therefore it can be written as (4.16)

$$\Omega(\theta) = \sum_k \sum_i \sum_j (W_{i,j}^{(k)})^2 \quad (4.15)$$

$$\Omega(\theta) = \sum_k \|W^{(k)}\|_F \quad (4.16)$$

The gradient to the regularizer with respect to the k^{th} layer is two times the weight matrix :

$$\nabla_{w^{(k)}} \Omega(\theta) = 2W^{(k)} \quad (4.17)$$

It is important to notice that this is applied only at the weights, because we don't expect to overfit the training set by changing the biases a lot but more by changing the weights that really determine how the function can become more or less non-linear.

4.3.2.2 Dropout Regularization

In deep neural networks have been proposed two ways of dealing with over-fitting , the first one is the unsupervised pre-training [26](which we will not discuss because it will not be used in this experiment therefore it is beyond the scope of this paper), the second is dropout [37]. Dropout was proposed by Geoffrey Hinton as a technique for performing regularization. It works by randomly dropping nodes in a neural network and it emulates ensembles¹ of neural networks. The application can be seen

¹train a group of prediction models, then average their prediction or take the majority vote

at figures 4.3, 4.4 where Figure 4.3 illustrates a standard neural network while Figure 4.4 illustrates the very same neural network after having applied dropout. In practice, for each hidden unit once it have been computed we will independently set it to 0 (dropping out the value derived from the training) with a probability ,usually of 0.5. This process continues till we reach the output layer. Therefore as a result of the whole process the hidden units can't collaborate with each other in order to generate complex patterns that might be useful to fit the training data so they are forced to extract a feature that is useful in general.

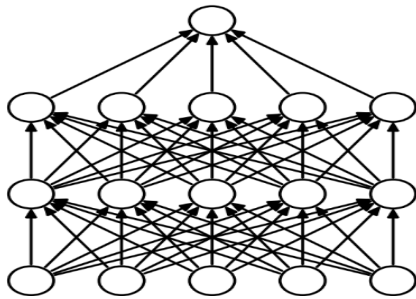


Figure 4.3: Standard Feed-forward Neural Network

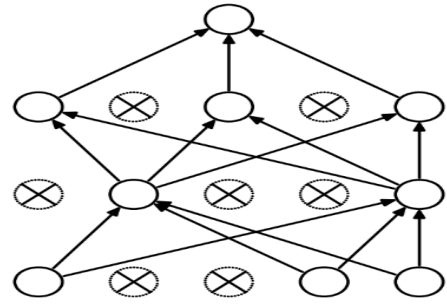


Figure 4.4: After applying Dropout

Source: <http://cs231n.github.io/neural-networks-2/>

The dropout regularization technique, it has an impact on both the forward and backward propagation algorithm for training a neural network. In more detail, regarding the forward propagation, we set a random binary mask $m^{(k)}$ with values 0 (dropout the weight of the unit) or 1 (retains the value), when it comes to the backward propagation, when we backpropagate the gradient till the preactivation(z1.3) we also need to multiply it by the mask vector, due to the chain rule. This practically means that many gradients will be set to zero so the backpropagation won't flow through the hidden units that were dropped out.

EXPERIMENTS

This chapter is dedicated to the practical comparison between the Constituent LSTM, with the Tree-Based GRU or Constituent GRU. It is important to notice that the experiments have been conducted 5 times and the results are the product of the averaged results of all the trials. We use the Stanford Sentiment Treebank(SST), and we use the standard train/dev/test splits of 6920/872/1821 for the binary classification subtask and 8544/1101/ 2210 for the fine-grained classification subtask (there are fewer examples for the binary subtask since the neutral instances have been excluded). The sentiment label at each node is predicted using the classifier covered at the next subsection 5.1.1. Moreover the SST have each sentence structured as constituent parse trees, therefore we will use the Constituent LSTM as a comparison to our model. Please find the code necessary for running those experiments to the next url: https://github.com/VasTsak/master_thesis.

5.1 Model comparison

Before proceeding to the experiments we made the assumption that the Tree-based GRU will be faster to be trained because of its fewer parameters. The experiments proved us right. But training speed is just a factor (not even that critical) to select a model, what we really care about is its capability of being able to identify the underlying pattern just right, not learn the training set by heart (overfitting) nor ignoring some important features(underfitting).

5.1.1 Classification model

The goal of the paper is to compare the performance of the Tree-GRU architecture against the Tree-LSTM architecture on sentiment classification tasks. In practice, the

model predicts a label $\hat{y}$ from a set of classes (2 for binary, 5 for fine grained) for some subset of nodes in a tree. The classifier and the objective function are exactly the same for both architectures. Let $\{x\}_j$ be the inputs observed at nodes in the subtree with root the node j .

$$\hat{p}_\theta(y | \{x\}_j) = \text{softmax}(W^p h_j + b^p), \quad (5.1)$$

$$\hat{y}_j = \underset{y}{\operatorname{argmax}} \hat{p}_\theta(y | \{x\}_j). \quad (5.2)$$

Let m be the number of labeled nodes in the training set and the superscript k be the k^{th} labeled node, the cost function is:

$$J(\theta) = -\frac{1}{m} \sum_{k=1}^m \log \hat{p}_\theta(y^{(k)} | \{x\}^{(k)}) + \frac{\lambda}{2} \|\theta\|_2^2 \quad (5.3)$$

5.1.2 Binary Classification

The binary classification is a problem that classifies whether the sentiment of the sentence is positive or negative. The process of the training can be seen at the Figures 5.1 , 5.2, and 5.3 where at Figure 5.1 it is plotted the average training time of each iteration and at Figures 5.2 are plotted the average loss of each iteration and at Figure 5.3 the training process¹ of Tree-GRU and Tree-LSTM.

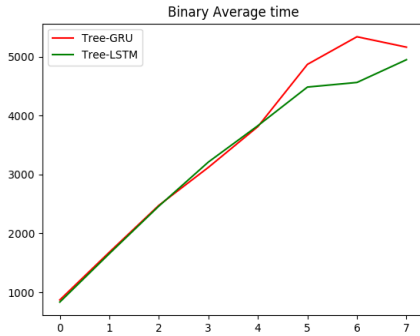


Figure 5.1: Binary Classification Average Training Time

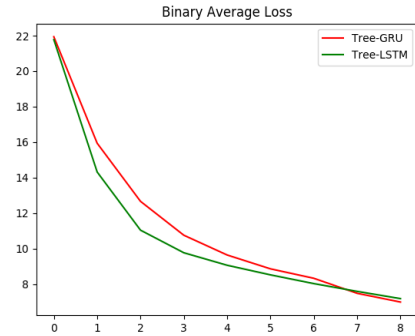


Figure 5.2: Binary Classification Average Training Loss

All the plots have as their x-axis the number of epochs. The metrics that we plot are computed till the 12th epoch and in cases of early stopping² we wouldn't take into account the 0 or "Non Assigned Number" of the trial that its training stopped earlier but we would just skip it and calculate the results based on the rest trials that had a full training process.

¹training process: the training and validation scores at each epoch.

²early stopping: when the validation error increases for a specified number of iterations, the training process stops

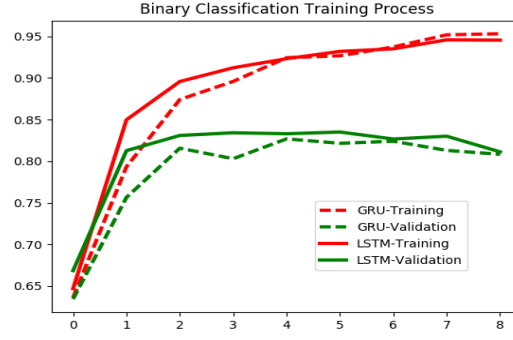


Figure 5.3: Binary Classification Training Process

5.1.3 Fine-grained Classification

The Fine-grained classification is a 5-class sentiment classification(1-Very Negative,2-Negative,3-Neutral,4-Positive,5-Very Positive). The experiment for the Fine-grained classification is under the same circumstances. The average time that each iteration lasted during the training process can be illustrated at Figure 5.4, the average loss that occurred during the training process is illustrated at Figure 5.5 and the training process of the fine grained classification can be seen at Figure 5.6.

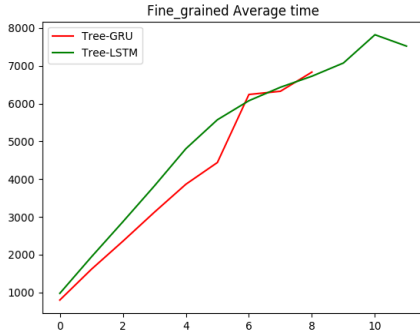


Figure 5.4: Fine Grained Classification Average Training Time

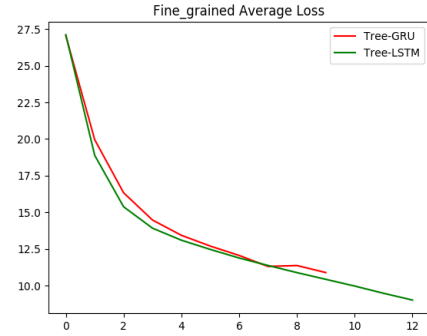
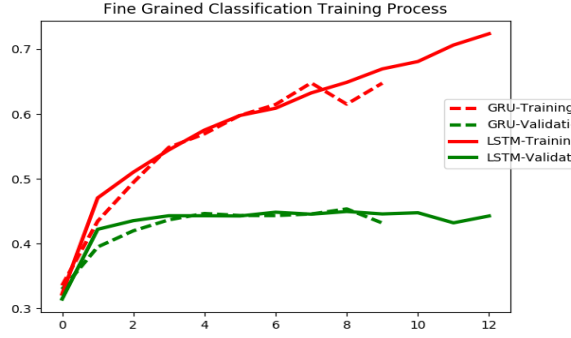


Figure 5.5: Fine Grained Classification Average Loss

5.1.4 Hyperparameters and Training Details

We have initialized the word representations using the pre-trained 300-dimensional GloVe vectors[60].The training of the model was done with AdaGrad[23] and a learning rate of 0.05 also we used the mini-batch gradient descent algorithm with batch size of 25. The model parameters were regularized with L2 regularization strength of 0.0001 and dropout rate of 0.5. For the training process we have applied the early stopping technique in order to avoid overfitting. The goal of this paper is not to achieve a state-of-art accuracy but to make a critical comparison between the two models therefore we

Figure 5.6: Fine Grained Classification Training Process



won't update the word representations during the training which boosts the accuracy approximately 0.05 (that is the accuracy boost gave to the Tree-LSTM).

5.2 Results

The results of both the binary and fine grained classification can be seen in Table 5.1 we can see that Tree-based GRU have slightly better performance with the tree-based LSTM, but it is important to notice from Table 5.2 the standard deviation of the individual predictions that the prediction from Tree-GRU seem to be more fluctuate than the ones from Tree-LSTM therefore it is possible that this difference of performance can be random, because the standard deviation is quite high for the case of Tree-GRU.

Something important to notice about the training process of fine grained classification is that, the Tree-based GRU would stop at the 8th iteration while the LSTM would go all the way till the 12th iteration. Moreover something else to notice is that the Tree-LSTM for fine-grained classification seems like it has some more training to do before it overfits, in contrast with the Tree-GRU which would overfit before having executed twelve iterations, which can be observed above (Figures 5.6). This may have to do with the hyperparameters that we have chosen. We have set the early stopping at 2 iterations(as the Tree-based LSTM paper had), if we would set it to 3 the Tree-GRU may keep on training till the 12th iteration.

Moreover Tree-GRU's training and validation scores seem to fluctuate more in the fine-grained classification 5.6 which may underlies unstable prediction and the need to train more.

Table 5.1: Sentiment Classification Accuracy

Model	Binary	Fine-grained
Tree-LSTM	84.43	45.71
Tree-GRU	85.61	46.43

Table 5.2: Sentiment Classification Standard Deviation

Model	Binary	Fine-grained
Tree-LSTM	0.35	0.55
Tree-GRU	0.93	0.98

5.3 Conclusions and Future Work

We can conclude that there is a difference in terms of performance ,not that significant though, between the tree-based LSTM and tree-based GRU. Moreover, Tree-based GRUs are trained faster -computationally- and Tree-based GRUs seem to converge faster so, the training process can stop earlier. Therefore it is a good alternative, if not a substitute. The area of Natural Language Processing is very active area of research, tree-based architectures proved to be very powerful for Natural Language Processing tasks, mostly because of their capability of handling negations. Many potential projects can be developed around Tree-Based GRUs, namely a Child-Sum approach, or the use of unique reset and update gate for each child, or even try different GRU architectures [22]).

For the end, one philosophical thought. Can you imagine an entire system of neural networks performing different tasks, so that the end result is something actionable. Like building a brain, the language processing system would be just a small part, but you may have a neural network to do part of speech tagging another neural network to do name entity recognition and another network to parse sentences into trees. Even a more challenging problem might be to figure out what is the general architecture we can use so that we don't even have to tell the system to learn these things. In other words, a network of neural networks where each neural network can figure out what it should do on its own and be useful for the overall system in a global manner.

BIBLIOGRAPHY

- [1] D. Bahdanau, K. Cho, and Y. Bengio. “Neural Machine Translation by Jointly Learning to Align and Translate.” In: *CoRR* abs/1409.0473 (2014). URL: <http://arxiv.org/abs/1409.0473>.
- [2] M. Baroni, G. Dinu, and G. Kruszewski. “Don’t count , predict ! A systematic comparison of context-counting vs . context-predicting semantic vectors.” In: *Proceedings of the 52nd Annual Meeting of the Association for Computational Linguistics*. (2014), pp. 238–247. ISSN: 1529-1898. DOI: [10.3115/v1/P14-1023](https://doi.org/10.3115/v1/P14-1023). arXiv: [arXiv:1011.1669v3](https://arxiv.org/abs/1011.1669v3).
- [3] Y Bengio. *Learning long-term dependencies with gradient descent is difficult*. 1994. DOI: [10.1109/72.279181](https://doi.org/10.1109/72.279181). arXiv: [arXiv:1211.5063v2](https://arxiv.org/abs/1211.5063v2).
- [4] Y. Bengio, R. Ducharme, P. Vincent, and C. Janvin. “A Neural Probabilistic Language Model.” In: *The Journal of Machine Learning Research* 3 (2003), pp. 1137–1155. ISSN: 15324435. DOI: [10.1162/153244303322533223](https://doi.org/10.1162/153244303322533223). arXiv: [arXiv:1301.3781v3](https://arxiv.org/abs/1301.3781v3).
- [5] Y. Bengio, N. Boulanger-Lewandowski, and R. Pascanu. “Advances in Optimizing Recurrent Networks.” In: *CoRR* abs/1212.0901 (2012). URL: <http://arxiv.org/abs/1212.0901>.
- [6] C. M. Bishop. “Neural Networks for Pattern Recognition.” In: (1995).
- [7] L. Bottou. “Large-Scale Machine Learning with Stochastic Gradient Descent.” In: *Proceedings of COMPSTAT’2010* (2010), pp. 177–186. ISSN: 0269-2155. DOI: [10.1007/978-3-7908-2604-3_16](https://doi.org/10.1007/978-3-7908-2604-3_16).
- [8] L. Bottou. “Stochastic Gradient Tricks.” In: *Neural Networks, Tricks of the Trade, Reloaded*. Vol. 7700. Springer, 2012, 430–445. URL: <https://www.microsoft.com/en-us/research/publication/stochastic-gradient-tricks/>.
- [9] C. Burgess and K. Lund. “The dynamics of meaning in memory.” In: *In Cognitive Dynamics: Conceptual Change in Humans and Machines*. Dietrich & Markman (Eds.), Psychology Press (2000), pp. 117–156. DOI: [10.1016/j.cognition.2007.07.015](https://doi.org/10.1016/j.cognition.2007.07.015). URL: papers3://publication/uuid/A6D7A23A-974A-4B5A-B501-8EA77A9306F2.

- [10] J. Chung, C. Gulcehre, K. Cho, and Y. Bengio. "Empirical Evaluation of Gated Recurrent Neural Networks on Sequence Modeling." In: *arXiv* (2014), pp. 1–9. arXiv: [1412.3555v1](#).
- [11] J. Chung, C. Gulcehre, K. Cho, and Y. Bengio. "Gated feedback recurrent neural networks." In: *Proceedings of the 32nd International Conference on Machine Learning, {ICML} 2015* 37 (2015), pp. 2067–2075. ISSN: 18792782. DOI: [10.1145/2661829.2661935](#). arXiv: [1502.02367](#). URL: <http://arxiv.org/abs/1502.02367>.
- [12] A. Coates, B. Carpenter, C. Case, S. Satheesh, B. Suresh, T. Wang, D. J. Wu, and A. Y. Ng. "Text Detection and Character Recognition in Scene Images with Unsupervised Feature Learning." In: *2011 International Conference on Document Analysis and Recognition*. 2011, pp. 440–445. DOI: [10.1109/ICDAR.2011.95](#).
- [13] R. Collobert. "Rehabilitation of Count-based Models for Word Vector Representations." In: (2014). arXiv: [arXiv:1412.4930v2](#).
- [14] R. Collobert and J. Weston. "A unified architecture for natural language processing." In: *Proceedings of the 25th international conference on Machine learning - ICML '08* 20.1 (2008), pp. 160–167. ISSN: 07224028. DOI: [10.1145/1390156.1390177](#). URL: <http://portal.acm.org/citation.cfm?id=1390177>{\%}5Cn<http://portal.acm.org/citation.cfm?doid=1390156>.1390177.
- [15] R. Collobert, J. Weston, L. Bottou, M. Karlen, K. Kavukcuoglu, and P. Kuksa. "Natural Language Processing (Almost) from Scratch." In: *Journal of Machine Learning Research* 12 (2011), pp. 2493–2537. ISSN: 0891-2017. DOI: [10.1.1.231.4614](#). arXiv: [1103.0398](#).
- [16] F Costa, P. Frasconi, V Lombardo, and G Soda. "Towards incremental parsing of natural language using recursive neural networks." In: *Applied Intelligence* 19.1-2 (2003), pp. 9–25. URL: <http://www.springerlink.com/index/WN8Q665MQ3462760.pdf>.
- [17] M. a. Covington. "A Fundamental Algorithm for Dependency Parsing." In: *Proceedings of the 39th Annual ACM Southeast Conference* (2001), pp. 95–102.
- [18] Y. Dauphin, R. Pascanu, C. Gulcehre, K. Cho, S. Ganguli, and Y. Bengio. "Identifying and attacking the saddle point problem in high-dimensional non-convex optimization." In: *arXiv* (2014), pp. 1–14. ISSN: 10495258. arXiv: [1406.2572](#). URL: <http://arxiv.org/abs/1406.2572>.
- [19] M.-C. De Marneffe, B. MacCartney, and C. D. Manning. "Generating typed dependency parses from phrase structure parses." In: *Proceedings of the 5th International Conference on Language Resources and Evaluation (LREC 2006)* (2006), pp. 449–454. DOI: [10.1.1.74.3875](#). URL: http://nlp.stanford.edu/pubs/LREC06{_}dependencies.pdf.

- [20] S. Deerwester, S. T. Dumais, G. W. Furnas, and T. K. Landauer. "Indexing by Latent Semantic Analysis." In: *Society* 41.6 (1990), pp. 391–407.
- [21] S. Diego, L. Papers, and H. Rohde. "Rhetorical questions as redundant interrogatives." In: *San Diego Linguistics Papers* 2.2 (2006).
- [22] A. Dosovitskiy, J. T. Springenberg, and T. Brox. "Learning to generate chairs with convolutional neural networks." In: *Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition* 07-12-June (2015), pp. 1538–1546. ISSN: 10636919. DOI: [10.1109/CVPR.2015.7298761](https://doi.org/10.1109/CVPR.2015.7298761). arXiv: [1512.03385](https://arxiv.org/abs/1512.03385).
- [23] J. Duchi. "Adaptive Subgradient Methods for Online Learning and Stochastic Optimization." In: 12 (2011), pp. 2121–2159.
- [24] J. Elith, C. H. Graham, R. P. Anderson, M. Dudík, S. Ferrier, A. Guisan, R. J. Hijmans, F. Huettmann, J. R. Leathwick, A. Lehmann, J. Li, L. G. Lohmann, B. A. Loiselle, G. Manion, C. Moritz, M. Nakamura, Y. Nakazawa, J. McC. M. Overton, A. Townsend Peterson, S. J. Phillips, K. Richardson, R. Scachetti-Pereira, R. E. Schapire, J. Soberón, S. Williams, M. S. Wisz, and N. E. Zimmermann. "Novel methods improve prediction of species' distributions from occurrence data." In: *Ecography* 29.2 (2006), pp. 129–151. ISSN: 09067590. DOI: [10.1111/j.2006.0906-7590.04596.x](https://doi.org/10.1111/j.2006.0906-7590.04596.x).
- [25] J. L. Elman. "Finding structure in time." In: *Cognitive Science* 14.2 (1990), pp. 179–211. DOI: [10.1016/0364-0213\(90\)90002-E](https://doi.org/10.1016/0364-0213(90)90002-E). URL: <http://groups.lis.illinois.edu/amag/langev/paper/elman90findingStructure.html>.
- [26] D. Erhan, A. Courville, and P. Vincent. "Why Does Unsupervised Pre-training Help Deep Learning ?" In: *Journal of Machine Learning Research* 11 (2010), pp. 625–660. ISSN: 15324435. DOI: [10.1145/1756006.1756025](https://doi.org/10.1145/1756006.1756025). arXiv: [arXiv: 1206.5538v1](https://arxiv.org/abs/1206.5538v1). URL: <http://portal.acm.org/citation.cfm?id=1756025>.
- [27] J. Firth. *Papers in linguistics, 1934-1951*. Oxford University Press, 1957. URL: <https://books.google.pt/books?id=yxZZAAAAMAAJ>.
- [28] F. Gers and J. Schmidhuber. "Recurrent nets that time and count." In: *Proceedings of the IEEE-INNS-ENNS International Joint Conference on Neural Networks. IJCNN 2000. Neural Computing: New Challenges and Perspectives for the New Millennium* 1 (2000), 189–194 vol.3. ISSN: 1098-6596. DOI: [10.1109/IJCNN.2000.861302](https://doi.org/10.1109/IJCNN.2000.861302). arXiv: [arXiv: 1011.1669v3](https://arxiv.org/abs/1011.1669v3). URL: <http://ieeexplore.ieee.org/lpdocs/epic03/wrapper.htm?arnumber=861302>.
- [29] F. a. Gers, N. N. Schraudolph, and J. Schmidhuber. "Learning Precise Timing with LSTM Recurrent Networks." In: *Journal of Machine Learning Research* 3.1 (2002), pp. 115–143. ISSN: 15324435. DOI: [10.1162/153244303768966139](https://doi.org/10.1162/153244303768966139). URL: http://www.crossref.org/jmlr{_}DOI.html.

- [30] C. Goller and A. Kuchler. “Learning task-dependent distributed representations by backpropagation through structure.” In: *Neural Networks, 1996., IEEE International Conference on*. Vol. 1. 1996, 347–352 vol.1. DOI: [10.1109/ICNN.1996.548916](https://doi.org/10.1109/ICNN.1996.548916).
- [31] G. H. Golub and C. Reinsch. “Singular value decomposition and least squares solutions.” In: *Numerische Mathematik* 14.5 (1970), pp. 403–420. ISSN: 0029599X. DOI: [10.1007/BF02163027](https://doi.org/10.1007/BF02163027).
- [32] G. H. Golub and C. F. Van Loan. *Matrix Computations*. 1996. DOI: [10.1063/1.3060478](https://doi.org/10.1063/1.3060478). arXiv: [arXiv:1011.1669v3](https://arxiv.org/abs/1011.1669v3).
- [33] A. Graves, A.-r. Mohamed, and G. Hinton. “Speech Recognition With Deep Recurrent Neural Networks.” In: *Icassp* 3 (2013), pp. 6645–6649. ISSN: 1520-6149. DOI: [10.1109/ICASSP.2013.6638947](https://doi.org/10.1109/ICASSP.2013.6638947). arXiv: [arXiv:1303.5778v1](https://arxiv.org/abs/1303.5778v1).
- [34] A. Gruslys, R. Munos, I. Danihelka, M. Lanctot, and A. Graves. “Memory-Efficient Backpropagation Through Time.” In: *CoRR* abs/1606.03401 (2016). URL: <http://arxiv.org/abs/1606.03401>.
- [35] J. Henderson. “Neural network probability estimation for broad coverage parsing.” In: *Proceedings of the tenth conference on European chapter of the Association for Computational Linguistics - EACL '03* 1.March (2003), pp. 131–138. DOI: [10.3115/1067807.1067826](https://doi.org/10.3115/1067807.1067826). URL: <http://portal.acm.org/citation.cfm?doid=1067807.1067826>.
- [36] G Hinton. “Mapping Part-Whole Hierachies into Connectionist Networks.” In: *Connectionist Symbol Processing* 46.1990 (1990), pp. 47–75.
- [37] G. E. Hinton, N. Srivastava, A. Krizhevsky, I. Sutskever, and R. Salakhutdinov. “Improving neural networks by preventing co-adaptation of feature detectors.” In: *CoRR* abs/1207.0580 (2012). URL: <http://arxiv.org/abs/1207.0580>.
- [38] S. Hochreiter. “Gradient Flow in Recurrent Nets: The Difficulty of Learning LongTerm Dependencies.” In: (2001). DOI: [10.1109/9780470544037.ch14](https://doi.org/10.1109/9780470544037.ch14). URL: <http://ieeexplore.ieee.org/search/srchabstract.jsp?arnumber=5264952>.
- [39] S. Hochreiter and J. Schmidhuber. “Long Short-Term Memory.” In: *Neural Comput.* 9.8 (Nov. 1997), pp. 1735–1780. ISSN: 0899-7667. DOI: [10.1162/neco.1997.9.8.1735](https://doi.org/10.1162/neco.1997.9.8.1735). URL: <http://dx.doi.org/10.1162/neco.1997.9.8.1735>.
- [40] A. E. Hoerl and R. W. Kennard. “Ridge Regression: Application to nonorthogonal problems.” In: *Technometrics* 12.1 (1970), pp. 69–82. ISSN: 0040-1706. DOI: [10.1080/00401706.1970.10488634](https://doi.org/10.1080/00401706.1970.10488634).
- [41] J. J. Hopfield. *Neural Networks and Physical Systems with Emergent Collective Computational Abilities*. 1982.

- [42] H. Jaeger. “The “ echo state ” approach to analysing and training recurrent neural networks – with an Erratum note 1.” In: *GMD Report* 148 (2010), pp. 1–47. ISSN: 18735223. DOI: [citeulike-article-id:9635932](#).
- [43] H. Jaeger, W. Maass, and J. Principe. *Special issue on echo state networks and liquid state machines*. 2007.
- [44] M. I. Jordan. *Serial order: A parallel distributed processing approach*. 1986. URL: <http://linkinghub.elsevier.com/retrieve/pii/S0166411597801112>{\%}5Cn<http://www.sciencedirect.com/science/article/pii/S0166411597801112>.
- [45] D. Klein and C. Manning. “Corpus-based induction of syntactic structure: models of dependency and constituency.” In: *ACL ’04: Proceedings of the 42nd Annual Meeting on Association for Computational Linguistics* 1 (2004), pp. 478–485. DOI: [10.3115/1218955.1219016](#). URL: <http://portal.acm.org/citation.cfm?id=1218955.1219016>{\%}5Cnpapers2://publication/uuid/3CEA60B9-E382-45DA-A4B5-F2771DFFC591.
- [46] D. Klein and C. D. Manning. “A Generative Constituent-Context Model for Improved Grammar Induction.” In: *Proceedings of the 40th Annual Meeting of the Association for Computational Linguistics (ACL)* July (2002), pp. 128–135. DOI: [10.3115/1073083.1073106](#).
- [47] K. J. Lang. “A Time Delay Neural Network Architecture for Speech Recognition.” AAI9011852. Doctoral dissertation. Pittsburgh, PA, USA, 1989.
- [48] O. Levy, Y. Goldberg, and I. Dagan. “Improving Distributional Similarity with Lessons Learned from Word Embeddings.” In: *Transactions of the Association for Computational Linguistics* 3 (2015), pp. 211–225. ISSN: 2307-387X. DOI: [10.1186/1472-6947-15-S2-S2](#). arXiv: [1103.0398](#). URL: <https://tacl2013.cs.columbia.edu/ojs/index.php/tacl/article/view/570>.
- [49] M. Li, T. Zhang, Y. Chen, and A. J. Smola. “Efficient Mini-batch Training for Stochastic Optimization.” In: *Kdd* (2014), pp. 661–670. ISSN: 03029743. DOI: [10.1145/2623330.2623612](#). arXiv: [1206.5533](#).
- [50] Z. C. Lipton. “A Critical Review of Recurrent Neural Networks for Sequence Learning.” In: *CoRR* abs/1506.0 (2015), pp. 1–38. ISSN: 9781450330633. DOI: [10.1145/2647868.2654889](#). arXiv: [1506.00019v2](#). URL: <http://arxiv.org/abs/1506.00019>.
- [51] J. Martens. “Deep learning via Hessian-free optimization.” In: *Proceedings of the 27th International Conference on Machine Learning (ICML-10)* 951 (2010), pp. 735–742. ISSN: 20901283. DOI: [10.1155/2011/176802](#). URL: [http://www.cs.toronto.edu/{~}asamir/cifar/HFO{_\}James.pdf](http://www.cs.toronto.edu/~asamir/cifar/HFO{_\}James.pdf).

- [52] W. S. McCulloch and W. Pitts. "A logical calculus nervous activity." In: *Bulletin of Mathematical Biology* 52.1 (1990), pp. 99–115. ISSN: 00074985. DOI: [10.1007/BF02478259](https://doi.org/10.1007/BF02478259).
- [53] T. Mikolov, K. Chen, G. Corrado, and J. Dean. "Distributed Representations of Words and Phrases and their Compositionality." In: *Nips* (2013), pp. 1–9. ISSN: 10495258. DOI: [10.1162/jmlr.2003.3.4-5.951](https://doi.org/10.1162/jmlr.2003.3.4-5.951). arXiv: [1310.4546](https://arxiv.org/abs/1310.4546).
- [54] T. Mikolov, K. Chen, G. Corrado, and J. Dean. "Efficient Estimation of Word Representations in Vector Space." In: *CoRR abs/1301.3781* (2013). URL: <http://arxiv.org/abs/1301.3781>.
- [55] A. Mnih. "Learning word embeddings efficiently with noise-contrastive estimation." In: *Nips* (2013), pp. 1–9. ISSN: 10495258. DOI: [10.3115/v1/P14-1023](https://doi.org/10.3115/v1/P14-1023). arXiv: [arXiv:1011.1669v3](https://arxiv.org/abs/1011.1669v3).
- [56] Y. Nesterov. "A method of solving a convex programming problem with convergence rate $O(1/k^2)$." In: *Soviet Mathematics Doklady*. Vol. 27. 2. 1983, pp. 372–376.
- [57] A. Ng. "1. Supervised learning." In: *Machine Learning* (2012), pp. 1–30.
- [58] C. Noam. "Syntactic Structures." In: (1958). DOI: [10.1515/9783110218329](https://doi.org/10.1515/9783110218329). URL: <http://www.degruyter.com/view/product/41408>.
- [59] B. Pang, L. Lee, and S. Vaithyanathan. "Thumbs up?: sentiment classification using machine learning techniques." In: *Proceedings of the Conference on Empirical Methods in Natural Language Processing* (2002), pp. 79–86. ISSN: 1554-0669. DOI: [10.3115/1118693.1118704](https://doi.org/10.3115/1118693.1118704). arXiv: [0205070 \[cs\]](https://arxiv.org/abs/0205070). URL: <http://portal.acm.org/citation.cfm?id=1118693.1118704>.
- [60] J. Pennington, R. Socher, and C. D. Manning. "GloVe : Global Vectors for Word Representation." In: (2014), pp. 1532–1543.
- [61] J. B. Pollack. "Recursive Distributed Representations." In: *Artif. Intell.* 46.1-2 (Nov. 1990), pp. 77–105. ISSN: 0004-3702. DOI: [10.1016/0004-3702\(90\)90005-K](https://doi.org/10.1016/0004-3702(90)90005-K). URL: [http://dx.doi.org/10.1016/0004-3702\(90\)90005-K](http://dx.doi.org/10.1016/0004-3702(90)90005-K).
- [62] N. Qian. "On the momentum term in gradient descent learning algorithms." In: *Neural Networks* 12.1 (1999), pp. 145–151. ISSN: 0893-6080. DOI: [http://dx.doi.org/10.1016/S0893-6080\(98\)00116-6](https://doi.org/10.1016/S0893-6080(98)00116-6). URL: <http://www.sciencedirect.com/science/article/pii/S0893608098001166>.
- [63] H. Ritter and T. Kohonen. "Self-organizing semantic maps." In: *Biological Cybernetics* 61.4 (1989), pp. 241–254. ISSN: 03401200. DOI: [10.1007/BF00203171](https://doi.org/10.1007/BF00203171).
- [64] F. Robinson A. J. Fallside. "The utility driven dynamic error propagation network." In: (1987).

- [65] F. Rosenblatt. “Principles of Neurodynamics. Perceptrons and the Theory of Brain Mechanisms.” In: *Archives of General Psychiatry* 7 (1962), pp. 218–219. ISSN: 0003-990X. DOI: [10.1001/archpsyc.1962.01720030064010](https://doi.org/10.1001/archpsyc.1962.01720030064010).
- [66] D. E. Rumelhart, G. E. Hinton, and R. J. Williams. *Learning representatons by back-propagating errors*. 1986. DOI: [10.1038/323533a0](https://doi.org/10.1038/323533a0). arXiv: [arXiv: 1011.1669v3](https://arxiv.org/abs/1011.1669v3).
- [67] G. Salton, A. Wong, and C. S. Yang. “A Vector Space Model for Automatic Indexing.” In: *Commun. ACM* 18.11 (Nov. 1975), pp. 613–620. ISSN: 0001-0782. DOI: [10.1145/361219.361220](https://doi.org/10.1145/361219.361220). URL: <http://doi.acm.org/10.1145/361219.361220>.
- [68] J. J. Schmidhuber. “Learning Complex, Extended Sequences Using the Principle of History Compression.” In: *Neural Comput.* 4.2 (1992), pp. 234–242. ISSN: 0899-7667. DOI: [10.1162/neco.1992.4.2.234](https://doi.org/10.1162/neco.1992.4.2.234). arXiv: [1103.0398](https://arxiv.org/abs/1103.0398). URL: <http://dx.doi.org/10.1162/neco.1992.4.2.234>.
- [69] R. Socher. “Recursive Deep Learning for Natural Language Processing and Computer Vision.” In: *PhD thesis* August (2014).
- [70] R. Socher, C. D. C. Manning, and A. Y. A. Ng. “Learning continuous phrase representations and syntactic parsing with recursive neural networks.” In: *Proceedings of the NIPS-2010 Deep Learning and Unsupervised Feature Learning Workshop* (2010), pp. 1–9. ISSN: 0302-9743. DOI: [10.1007/978-3-540-87479-9](https://doi.org/10.1007/978-3-540-87479-9). URL: <http://wuawua.googlecode.com/files/LearningContinuousPhraseRepresentationsandSynta.pdf>.
- [71] R. Socher, E. Huang, and J. Pennington. “Dynamic Pooling and Unfolding Recursive Autoencoders for Paraphrase Detection.” In: *Advances in Neural Information Processing Systems* (2011), pp. 801–809. ISSN: 9781618395993. URL: http://machinelearning.wustl.edu/mlpapers/paper/{_}files/NIPS2011/{_}0538.pdf{\\}%5Cnhttps://papers.nips.cc/paper/4204-dynamic-pooling-and-unfolding-recursive-autoencoders-for-paraphrase-detection.pdf.
- [72] R. Socher, B. Huval, C. D. Manning, and A. Y. Ng. “Semantic Compositionality through Recursive Matrix-Vector Spaces.” In: *Proceedings of the 2012 Joint Conference on Empirical Methods in Natural Language Processing and Computational Natural Language Learning* Mv (2012), pp. 1201–1211. ISSN: 9781937284435. DOI: [10.1162/153244303322533223](https://doi.org/10.1162/153244303322533223). arXiv: [arXiv: 1301.3781v3](https://arxiv.org/abs/1301.3781v3).
- [73] R. Socher, A. Perelygin, and J. Wu. “Recursive deep models for semantic compositionality over a sentiment treebank.” In: *Proceedings of the ...* (2013), pp. 1631–1642. ISSN: 1932-6203. DOI: [10.1371/journal.pone.0073791](https://doi.org/10.1371/journal.pone.0073791). arXiv: [1512.03385](https://arxiv.org/abs/1512.03385). URL: http://nlp.stanford.edu/{~}socherr/EMNLP2013/{_}RNTN.pdf{\\}%5Cnhttp://www.aclweb.org/anthology/D13-1170{\\}%5Cnhttp://

- aclweb.org/supplementals/D/D13/D13-1170.Attachment.pdf{\%}5Cnhttp://oldsite.aclweb.org/anthology-new/D/D13/D13-1170.pdf.
- [74] I. Sutskever. “Training Recurrent neural Networks.” In: *PhD thesis* (2013), p. 101. arXiv: 1456339 [arXiv:submit].
- [75] I. Sutskever, O. Vinyals, and Q. V. Le. “Sequence to Sequence Learning with Neural Networks.” In: CoRR abs/1409.3215 (2014). URL: <http://arxiv.org/abs/1409.3215>.
- [76] R. S. Sutton. “Learning to Predict by the Method of Temporal Differences.” In: *Machine Learning* 3.1 (1988), pp. 9–44. ISSN: 08856125. DOI: 10.1023/A:1018056104778. URL: citeseer.ist.psu.edu/sutton88learning.html.
- [77] K. S. Tai, R. Socher, and C. D. Manning. “Improved Semantic Representations From Tree-Structured Long Short-Term Memory Networks.” In: *Acl (1)* (2015), pp. 1556–1566. ISSN: 9781941643723. DOI: 10.1515/popets-2015-0023. arXiv: 1503.0075. URL: <http://aclweb.org/anthology/P/P15/>.
- [78] B. Taskar, D. Klein, M. Collins, D. Koller, and C. Manning. “Max-Margin Parsing.” In: *Proc. EMNLP* (2004), pp. 1–8.
- [79] O. Vinyals, A. Toshev, S. Bengio, and D. Erhan. “Show and Tell: A Neural Image Caption Generator.” In: CoRR abs/1411.4555 (2014). URL: <http://arxiv.org/abs/1411.4555>.
- [80] P. J. Werbos. “Backwards Differentiation in AD and Neural Nets: Past Links and New Opportunities.” In: *Lecture Notes in Computational Science and Engineering* 50 (2006), pp. 15–34. ISSN: 14397358. DOI: 10.1007/3-540-28438-9_2.
- [81] W. Zaremba and I. Sutskever. “Learning to Execute.” In: CoRR abs/1410.4615 (2014). URL: <http://arxiv.org/abs/1410.4615>.
- [82] M. D. Zeiler. “ADADELTA: An Adaptive Learning Rate Method.” In: CoRR abs/1212.5701 (2012). URL: <http://arxiv.org/abs/1212.5701>.
- [83] M. Zhu, Y. Zhang, W. Chen, M. Zhang, and J. Zhu. “Fast and Accurate Shift-Reduce Constituent Parsing.” In: *Proceedings of the 51st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)* (2013), pp. 434–443. URL: <http://www.aclweb.org/anthology/P/P13/P13-1043.pdf>{\%}5Cnhttp://www.aclweb.org/anthology/P13-1043.
- [84] D. Zipser and R. J. Williams. “Gradient-Based Learning Algorithms for Recurrent Networks and Their Computational Complexity.” In: *Back-propagation: Theory, Architectures and Applications* (1995), pp. 433–486.



APPENDIX 1

Let E is the error, $\delta h_t = \frac{\partial E}{\partial h_t}$, we seek to find δo_t , δC_t , δi_t , $\delta \tilde{C}_t$, δf_t , δC_{t-1} .

$$\delta o_t = \frac{\partial E}{\partial o_t} = \frac{\partial E}{\partial h_t} \frac{\partial h_t}{\partial o_t} = \delta h_t \odot \tanh(C_t) \quad (\text{A.1})$$

$$\delta C_t = \frac{\partial E}{\partial C_t} = \frac{\partial E}{\partial h_t} \frac{\partial h_t}{\partial C_t} = \delta h_t \odot o_t \odot (1 - \tanh^2(C_t)) \quad (\text{A.2})$$

$$\delta i_t = \frac{\partial E}{\partial i_t} = \frac{\partial E}{\partial C_t} \frac{\partial C_t}{\partial i_t} = \delta C_t \odot \tilde{C}_t \quad (\text{A.3})$$

$$\delta \tilde{C}_t = \frac{\partial E}{\partial \tilde{C}_t} = \frac{\partial E}{\partial C_t} \frac{\partial C_t}{\partial \tilde{C}_t} = \delta C_t \odot i_t \quad (\text{A.4})$$

$$\delta f_t = \frac{\partial E}{\partial f_t} = \frac{\partial E}{\partial C_t} \frac{\partial C_t}{\partial f_t} = \delta C_t \odot C_{t-1} \quad (\text{A.5})$$

$$\delta C_{t-1} = \frac{\partial E}{\partial C_{t-1}} = \frac{\partial E}{\partial C_t} \frac{\partial C_t}{\partial C_{t-1}} = \delta C_t \odot f_t \quad (\text{A.6})$$

A P P E N D I X

APPENDIX 2

Given $\delta h_t = \frac{\theta E}{\theta h_t}$ we seek to find $\delta \tilde{h}_t h_t$, δr_t and δz_t

$$\delta \tilde{h}_t = \frac{\theta E}{\theta \tilde{h}_t} = \frac{\theta E}{\theta h_t} \frac{\theta h_t}{\theta \tilde{h}_t} = \delta h_t (1 - z) \quad (\text{B.1})$$

$$\delta r_t = \frac{\theta E}{\theta r_t} = \frac{\theta E}{\theta \tilde{h}_t} \frac{\theta \tilde{h}_t}{\theta r_t} = \delta \tilde{h}_t \odot h_{t-1} U^h \odot (1 - \tilde{h}_t^2) \quad (\text{B.2})$$

$$\delta z_t = \frac{\theta E}{\theta z_t} = \frac{\theta E}{\theta h_t} \frac{\theta h_t}{\theta z_t} = \delta h_t (h_{t-1} - \tilde{h}_t) \quad (\text{B.3})$$

