

DOKUZ EYLÜL UNIVERSITY
GRADUATE SCHOOL OF NATURAL AND APPLIED SCIENCES

TR-SUM: A TEXT SUMMARIZER
FOR TURKISH

by
Yiğit YÜKSEL

December, 2021
İZMİR

TR-SUM: A TEXT SUMMARIZER FOR TURKISH

**A Thesis Submitted to the
Graduate School of Natural and Applied Sciences of Dokuz Eylül University
In Partial Fulfillment of the Requirements for the Degree of Master of Sciences
in Computer Engineering, Applied Computer Engineering**

**by
Yiğit YÜKSEL**

**December, 2021
İZMİR**

M.Sc THESIS EXAMINATION RESULT FORM

We have read the thesis entitled “**TR-SUM: A TEXT SUMMARIZER FOR TURKISH**” completed by **Yiğit YÜKSEL** under supervision of **PROF. DR. Yalçın ÇEBİ** and we certify that in our opinion it is fully adequate, in scope and in quality, as a thesis for the degree of Master of Science.

.....
Prof. Dr. Yalçın ÇEBİ

Supervisor

.....

Assoc Prof. Dr. Gökhan DALKILIÇ
(Jury Member)

.....

Assist. Prof Dr. Ömer AYDIN
(Jury Member)

Prof. Dr. Okan FISTIKOĞLU

Director

Graduate School of Natural and Applied Sciences

ACKNOWLEDGMENTS

Firstly, I would like to express my special respect and sincere thanks of gratitude and to my thesis advisor Prof. Dr. Yalçın Çebi for his valuable guidance during my thesis studies. It was a great privilege to complete this thesis with his guidance. Secondly, I would particularly thank Prof. Dr. Selda Akçalı and Asst. Prof. Tolga Çelik from Journalism Department, Faculty of Communication, Ege University, and the students of this department for their contribution during the dataset collection period of this thesis.

In addition, I would like to thank my beloved wife Damla Yüksel for her continuous support and encouragement in each step of my studies.

Yiğit YÜKSEL

TR-SUM: A TEXT SUMMARIZER FOR TURKISH

ABSTRACT

In today's world of big data, the tracing and obtaining of specific and useful information from stored and textual data is one of the most significant challenge of summarization. Summarization is basically the process of shortening the text. However, retrieving a short summary or only relevant information from large texts is demanding. The summary should include the main purpose, main idea, and important information of the text. The studies for both the extractive and abstractive text summarization models are recently increasing for various languages. However, fewer studies have been carried out for the abstractive text summarization of Turkish language. In addition, there are less collected datasets to be summarized for Turkish language in the literature. Thus, the contribution of this thesis is in two-fold. Firstly, a news dataset is collected for Turkish language. Secondly, this thesis proposes the adaptation and implementation of three abstractive deep neural network models for the Turkish language. These models are (i) Attention Based Seq2Seq Neural Network, (ii) Pointer Generator Seq2Seq Neural Network and (iii) Reinforcement Learning with Seq2Seq Neural Network. All three models are preprocessed with both ConceptNet-Numberbatch word embedding and fastText word embedding. Then, these models are evaluated on the collected Turkish dataset based on the ROUGE-1, ROUGE-2, and ROUGE-L scores. According to the computation experimentation, All ROUGE scores of each neural network model that is studied with fastText word embedding have decent quality. The highest ROUGE scores are acquired by Pointer Generator Seq2Seq Neural with fastText word embedding. This indicates that Pointer Generator Seq2Seq Neural Network is a promising deep neural network model that may produce well-qualified text summaries.

Keywords: Text summarization, abstractive text summarization, extractive text summarization, summarization for Turkish language, encoder-decoder architecture, sequence-to-sequence models, long short-term memory



TR-SUM: TÜRKÇE İÇİN BİR METİN ÖZETLEYİCİ

ÖZ

Günümüzün büyük veri dünyasında, saklanan ve metinsel verilerden belirli ve faydalı bilgilerin izlenmesi ve elde edilmesi, özetlemenin en önemli zorluklarından biridir. Özetleme, temel olarak metni kısaltma işlemidir. Ancak, büyük metinlerden kısa bir özet veya yalnızca ilgili bilgileri almak zordur. Özet, metnin ana amacını, ana fikrini ve önemli bilgilerini içermelidir. Son zamanlarda çeşitli diller için hem çıkarımsal hem de soyutlayıcı metin özetleme modellerine yönelik çalışmalar artmaktadır. Ancak Türkçe dili için soyutlayıcı metin özetlemesi alanında daha az sayıda çalışma yapılmıştır. Ayrıca literatürde Türkçe için özetlenecek daha az sayıda veri seti bulunmaktadır. Dolayısıyla bu tezin katkısı iki yönlüdür. İlk olarak Türkçe dili için bir veri seti toplanmıştır. İkinci olarak, bu çalışmada üç adet soyutlayıcı derin sinir ağı modelinin Türkçe diline uyarlanması ve uygulanması önerilmektedir. Bu modeller, (i) Dikkat Tabanlı Sıradan Sıraya Yapay Sinir Ağı, (ii) Pointer Generator Sıradan Sıraya Yapay Sinir Ağı ve (iii) Sıradan Sıraya Yapay Sinir Ağı ile Güçlendirmeli Öğrenme'dir. Her üç model de hem ConceptNet-Numberbatch kelime gömme hem de fastText kelime gömme ile önceden işlenmiştir. Daha sonra bu modeller, ROUGE-1, ROUGE-2 ve ROUGE-L puanlarına dayalı olarak Türkçe için toplanan veri seti üzerinde değerlendirilmiştir. Hesaplama deneyine göre, fastText kelime gömme ile çalışılan her bir sinir ağı modelinin tüm ROUGE puanları iyi kaliteye sahiptir. En yüksek ROUGE puanları, Pointer Generator Sıradan Sıraya Yapay Sinir Ağı tarafından fastText kelime gömme ile elde edilir. Bu, Pointer Generator Sıradan Sıraya Yapay Sinir Ağı'nın iyi nitelikli metin özetleri üretebilen umut verici bir derin sinir ağı modeli olduğunu gösterir.

Anahtar kelimeler: Metin özetleme, soyut metin özetleme, çıkarıcı metin özetleme, Türk Dili için özetleme, kodlayıcı-kodçözücü mimarisi, sıradan sıraya modeller, uzun kısa süreli bellek



CONTENTS

M.Sc THESIS EXAMINATION RESULT FORM	ii
ACKNOWLEDGMENTS.....	iii
ABSTRACT	iv
ÖZ.....	vi
LIST OF FIGURES.....	xi
LIST OF TABLES	xii
 CHAPTER ONE - INTRODUCTION	1
 1.1 General Overview	1
1.2 Motivation.....	1
1.3 Aim of Thesis.....	2
1.4 Contributions	2
1.5 Statement of Thesis.....	3
 CHAPTER TWO - RELATED WORKS.....	4
 2.1 Overview.....	4
2.2 Text Summarization Studies.....	4
2.2.1 Extractive Text Summarization Studies.....	5
2.2.2 Abstractive Text Summarization Studies.....	8

2.2.3	Text Summarization Studies for Turkish.....	13
2.3	Evaluation of Related Works.....	15
CHAPTER THREE - TEXT SUMMARIZATION DATASETS AND METHOD		17
3.1	Datasets.....	17
3.1.1	Datasets in English	17
3.1.2	Datasets in Turkish.....	19
3.2	Encoder-Decoder Sequence to Sequence (Seq2Seq) Model	20
CHAPTER FOUR - TR-SUM: A TEXT SUMMARIZER FOR TURKISH.....		23
4.1	General Overview of “TR-Sum: A Text Summarizer for Turkish”	23
4.2	TR-News-Sum Dataset	24
4.3	Data Pre-Processing.....	26
4.4	The Proposed Neural Network Models for Turkish Text Summarization.....	29
CHAPTER FIVE - APPLICATION AND EVALUATION OF TR-SUM: A TEXT SUMMARIZER FOR TURKISH.....		35
5.1	Application.....	35
5.2	Evaluation	36

CHAPTER SIX - CONCLUSION AND FUTURE WORK.....	42
--	-----------

REFERENCES.....	46
------------------------	-----------

APPENDICES	58
-------------------------	-----------



LIST OF FIGURES

Figure 3.1 General Seq2Seq Architecture for Text Translation.....	22
Figure 4.1 Flow Chart of TR-Sum: A Text Summarizer for Turkish	23
Figure 4.2 A Screenshot of the Data Collection System.....	25
Figure 4.3 Database Schema of the System	26
Figure 4.4 Flow Chart of Data Pre-Processing	27
Figure 4.5 Architecture of the Attention Based Seq2Seq Neural Network	30
Figure 4.6 Architecture of the Pointer Generator Seq2Seq Neural Network.....	32
Figure 4.7 Architecture of the Reinforcement Based Seq2Seq Neural Network	33
Figure 5.1 Pseudocode of the Models	35

LIST OF TABLES

Table 3.1 Datasets in English	17
Table 4.1 Stop Words	28
Table 4.2 The Characteristics of the “TR-News-Sum” Dataset	29
Table 4.3 An Example from Dataset as Input Vector	29
Table 5.1 Parameter Settings.....	36
Table 5.2 Average Rouge Scores of the Three Neural Network Models	39
Table 5.3 Improvement Achieved by <i>fastText</i> Word Embedding on Three Neural Network Models	40

CHAPTER ONE

INTRODUCTION

1.1 General Overview

In today's world of big data, the tracing and obtaining of specific and useful information from stored and textual data is one of the most significant challenge of summarization. Summarization is basically the process of shortening the text. However, retrieving a short summary or only relevant information from large texts is demanding. The summary should include the main target, main idea and important information in the text. Most of the diverse fields such as public opinion monitoring, news recommendation, content/article summarization etc. concentrate on the summarization (Fang, Mu, Deng, & Wu, 2017). There will be more demand on the summarization of long-sized texts and electronic documents since they growth an exponentially (Uçkan & Karcı, 2020). Thus, a need for automated text summarization tool is inevitable in the existence of long-sized texts and electronic documents.

1.2 Motivation

The number of electronic documents, online articles, especially electronic news, has grown exponentially over the last decade (Hark, Uckan, Seyyarer, & Karcı, 2019), and because of this abundance, sorting of information and retrieving a short summary from the long-sized texts has been considered a significant research area (Altan, 2004). Therefore, effective text summarization systems, are required in all languages. Effective text summarization systems must be as fast as possible in terms of summary production and be able to produce good-quality summaries in terms of the quality of the summaries. However, to be able to implement effective text summarization systems, a dataset constituting original texts and their summaries is a must. In other words, there should be a reference dataset where the summaries for the long-sized texts exist and are reachable so that the text summarization systems can be applied to the dataset. In this way, the text

summarization systems can generate their own text summaries, and hence the performance of these systems can be compared with the existed summaries in the dataset. By the performance of these systems, it has meant that the quality of the generated text summaries by the text summarization system.

There are two main motivations for this thesis. The first one is to generate a novel news dataset for the Turkish language. This generated news dataset aims to be a basic dataset for researchers to test further different algorithms on. The latter one is to summarize the long-sized news texts of the generated news dataset for the Turkish language with an effective text summarization system.

1.3 Aim of Thesis

Text summarization algorithms are required for today's big data environment in all languages. The need for summary retrieval from long-sized text is inevitable in today's big data. According to the research and review performed on the literature so far, an abstractive text summarization system for Turkish is a very significant research direction. This thesis fundamentally aims to generate an abstractive text summarizer for Turkish language depending on the gap in the literature of Turkish text summarization. Firstly, a web interface of the data collection system to collect data from users will be constituted. Then, a news dataset for Turkish language will be generated with the collected data. The collected dataset will be preprocessed to be fed to the deep neural network models. Following, abstractive deep neural network models will be studied on the generated dataset, and the performance of the models will be compared regarding the various ROUGE scores. At the end, an automatic abstractive text summarizer for Turkish language will be developed.

1.4 Contributions

Several contributions are achieved with this thesis. Firstly, the related works in the literature of extractive and abstractive text summarization are analyzed extensively. Then, the works for the Turkish language are also explored in depth. As the output of this

literature review, all studies are reported and categorized in terms of their methodology, used dataset, and evaluation metric information. Consequently, depending on the gap that is depicted by literature review, automatic abstractive text summarizer for Turkish language become the research direction of this thesis. A web interface of the data collection system is constituted to collect data from users. Then, a novel news dataset named as “TR-News-Sum” dataset is generated with the collected data. The collected dataset is preprocessed to be feed to the deep neural network models. Two different word embeddings, *ConceptNet-Numberbatch* and *fastText* word embeddings, are used during the data pre-processing. Following, three deep neural network models are studied on the “TR-News-Sum” dataset. These models are (i) Attention Based Seq2Seq Neural Networks, (ii) Pointer Generator Seq2Seq Neural Networks and (iii) Reinforcement Learning with Seq2Seq Neural Networks. All three models are preprocessed with both *ConceptNet-Numberbatch* word embedding and *fastText* word embeddings. The performance of the deep neural network models is measured on the collected Turkish dataset based on the ROUGE-1, ROUGE-2, and ROUGE-L scores. According to the computational experimentation, *Pointer Generator Seq2Seq Neural Network* is a promising deep neural network model that can procure well-qualified text summaries for the Turkish language.

1.5 Statement of Thesis

This thesis is structured as follows: In Chapter 2, firstly, the related works in the literature of extractive and abstractive summarization methods are presented. Then, the related works performed for Turkish language are provided. Later, in Chapter 3, text summarization data sets and method are discussed briefly. While the proposed text summarizer for Turkish (*TR-Sum: A Text Summarizer for Turkish*) is presented in Chapter 4, the application and evaluation of *TR-Sum: A Text Summarizer for Turkish* method are exhibited in Chapter 5. Finally, the concluding comments and possible future research directions are stated in Chapter 6.

CHAPTER TWO

RELATED WORKS

2.1 Overview

Recently, text summarization systems are a trending need of people who deals with a huge amount of digital text data worldwide (Yousefi-Azar & Hamey, 2017). Most of the people generally encounter an extensive amount of text datasets in their working life. Rather than reading the entire text data, reading its summary is usually preferable. Not only in working lives, but also in daily lives a need for text summarization arises, as well. Hence, in recent years, researchers more likely tend to cope with text summarization systems in all languages.

In this chapter, firstly, the related works in the literature of extractive and abstractive text summarization methods are provided in Sections 2.2.1 and 2.2.2, respectively. Then, the literature review of summarization methods in text summarization specifically in Turkish language are given in Section 2.2.3. Finally, the evaluation of related works as well as the gap in the literature are explicitly shown by concluding the motivation of this thesis in Section 2.3.

2.2 Text Summarization Studies

Text summarization can be classified into three main categories, based on input type, based on output type, and based on purpose. In this thesis, based on output type summarization is considered. Mainly, two different types of approaches are used based on output type:

- i. Extractive Text Summarization
- ii. Abstractive Text Summarization

Extractive text summarization approach selects important words/phrases from the source document and combines them into a summary. As opposed to extractive

summarization, abstractive text summarization approach generates new sentences which contain the essence and meaning of the source document.

2.2.1 Extractive Text Summarization Studies

As one of the early studies on extractive text summarization methods, Erkan & Radev (2004) proposed a model that computes the value of the sentences with the graph-based algorithms. The eigenvector centrality is used to determine value of the sentence and the intra-sentence cosine similarity is used to determine relations between the sentences. The outcome of the study demonstrates that the LexRank degree-based algorithm mostly dominates the other degree-based and centroid-based algorithms. As another early study, the model of Barzilay & McKeown (2005) extracts the frequent data from the similar sentences and generates the summarized text from the extracted data. The generation is handled by eliminating and paraphrasing the input data. Then, Amancio, Nunes, Oliveira, & Costa (2012) introduced a diversity metric that has been calculated with betweenness, vulnerability, and diversity in Brazilian Portuguese language. The summarization results of the proposed model are superior against the statistical models performed Brazilian Portuguese language. In terms of the multi document analysis, Ferreira et al. (2014) proposed a graph based multi document model that synthesizes the important information from the documents. The model also tries to find solutions on extractive summarization problems: redundancy, and diversity.

Later, Fang, Mu, Deng, & Wu (2017) presented a sentence scoring technique which combines the word-sentence level relations with the graph-based unsupervised model called as CoRank. Additionally, the quality of the summarization was supported with the redundancy elimination techniques. In another study, a model that combines two different models was presented by Jain, Bhatia, & Thakur (2017). The first model of Jain et al. (2017) extracts the features from the document whereas the second model is a supervised neural network model to generate the summaries. Nallapati, Zhai, & Zhou (2017) proposed a model that uses the abstractive training system to improve the extractive labelling training time. The model has also intuitive visualization feature that allows to

illustrate the decisions made by the model. Yousefi-Azar & Hamey (2017) generated a deep neural network model with encoders. The model calculates the term-frequency as a feature space to feed the neural network model. Furthermore, the study also investigated the impacts of adding random noises in feature space summarization results.

Then, in the study of Al-Sabahi, Zuping, & Nadher (2018), they mimicked the structure of the input document to be able to yield better results. Furthermore, the model also computes the probability of the sentence-summary relations to decide on which part it should pay attention to. Next, as an iterative process, Chen et al. (2018) introduced a model that passes over the document multiple times to reach better summarization results. Besides, the model contains a selective reading mechanism that decides on which part of the summarization should be updated over the iterations. Furthermore, in the study of Jadhav & Rajan (2018), the model detects the important sentences and the key words on the document then use them to create summaries. The model computes the relations between the sentences and the keywords with the two-level pointer network-based architecture. Also, Khatri, Singh, & Parikh (2018) proposed that the seq2seq models should be feed by contextual information before starting the encoding of the input. Through this approach, they aimed to have more relevant summarization results to the context.

As a multi-objective analysis, Sanchez-Gomez, Vega-Rodríguez, & Pérez (2018) presented a MOABC-based approach to optimize the extractive text summarization process. The optimization process focuses the following areas information extraction and removing redundant information. Sinha, Yadav, & Gahlot (2018) proposed a data-driven approach using neural networks in which the model divides the document into chunks in a fixed size, and then it feeds the network with the chunks. Wu & Hu (2018) presented a two-phase model. The first model detects the cross-sentence semantics and the syntactic coherence relations. Then, the output of the first model is given as input for the Reinforced Neural Extractive Summarization (RNES) model. The second model optimizes the coherence and prioritizes the sentences for the summary.

As of 2019, Brito, Lübbering, Biesner, Hillebrand, & Bauckhage (2019) carried out a recurrent neural-network-based model to decide on whether the sentence of the input text belongs to generated summary or not. Then, Keneshloo, Ramakrishnan, & Reddy (2019) applied transfer learning in their study. They presented a reinforcement neural network model enhanced with a self-critic policy gradient approach. Lierde & Chow (2019) proposed a graph-based text summarization model that treats the nodes as sentences and topics. The topics are extracted as relations between the sentences in which a probabilistic model was carried out. Liu & Lapata (2019) provided a model that is based on BERT (Devlin, Chang, Lee, & Toutanova, 2019) for the document summarization. The model can extract the semantics of a document and find the sentence references. Furthermore, with the help of the two-staged fine-tuning part, generated summary quality is increased. Moreover, a statistical model that ranks the sentences by their assigned weights was proposed by Madhuri & Kumar (2019). In this study, highly ranked sentences are selected to the summarized document. Xu & Durrett (2019) proposed an extraction model with a compression algorithm that decides on which part of the document should be removed. Lastly, Xu, Gan, Cheng, & Liu (2019) generated a graph-based BERT model. The graph part of the model aims to find the long-range sentence dependencies to reduce redundant words/sentences in the summary.

As of 2020, Bhargava & Sharma (2020) implemented paraphrase detection to conclude that whether a sentence is related to summary or not. The algorithm of Elbarougy, Behery, & El Khatib (2020) fundamentally consider the document as a graph and hence the sentences as the vertices. A new algorithm named as Modified Rank algorithm is implemented to the generated graphs. The leading study to integrate supervised and unsupervised learning algorithms with each other is the study of Ghodrattnama, Beheshti, Zakershahrak, & Sobhanmanesh (2020) where they generated an extractive model regarding a dynamic feature weighting. While a genetic algorithm based lexical semantic analysis is implemented by Hernandez-Castaneda, Garcia-Hernandez, Ledeneva, & Millan-Hernandez, 2020, a distributional semantic model for ranking is developed by Mohd, Jan, & Shah, 2020. Sanchez-Gomez, Vega-Rodríguez, & Pérez, 2020 is presented

a multi-objective artificial bee colony (MOABC) based extractive text summarization. A novel approach with maximum independent sets for extractive text summarization is proposed by Uçkan & Karıcı, 2020.

A summary table for the studies related to the extractive text summarization algorithms is given in Appendix A in chronological order.

2.2.2 Abstractive Text Summarization Studies

In 2015, the studies of Lopyrev (2015) and Rush, Chopra, & Weston (2015) are published. Lopyrev (2015) proposed an encoder-decoder RNN with an LSTM model that generates summaries using only the first 50 words of an input text. The model produces a succinct summary from the first fifty words which is mostly correct grammatically. The study states that simplified attention mechanism performs better than complex attention mechanisms. Rush et al. (2015) also presented an attention-based model incorporating with a probabilistic model which results with more accurate summaries. Since the model is simple, the training process can be readily adapted to a huge amount of training input. The proposed model of Rush et al. (2015) is open for grammatically improvement and adaptation for long texts such as paragraphs.

Chen, Zhu, Ling, Wei, & Jiang (2016) presented an attention-based model combined with a distraction model that leads to focusing on different parts of the source text. The distraction model enables the model to embrace the more general meaning of the source text. The output of the study shows that the distraction model helps to achieve novel results, especially on the long length documents. Followingly, an attention-based encoder-decoder model which has novel performances on two different corpora proposed in the study of Nallapati, Zhou, Santos, & Xiang (2016). The model resolves crucial concerns in summarization, such as modeling keywords, finding the relation between sentence and word, revealing words that unknown words during training. The study also contributes to the literature with a new dataset, to be used in the evaluation of other studies. Next, an attention-based encoder-decoder model that leverages structural syntactic and semantic information which is named as Abstract Meaning Representation (AMR) was proposed

by Takase, Suzuki, Okazaki, Hirao, & Nagata (2016). The study illustrates that Abstract Meaning Representation (AMR) helps to obtain improved results on the summarized texts. As new features, two additional instruments (Read Again & Copy Mechanism) those alter the problems of encoder-decoder models are introduced to the literature by Zeng, Luo, Fidler, & Urtasun (2016). “Read Again” model waits to process input words until the input sequence is processed completely. “Copy Mechanism” allows handling non-dictionary words which improves process time and decreases the vocabulary size.

As a user friendly study, Fan, Grangier, & Auli (2017) presented a model that enables users to decide summarization result attributes such as length, style, entities etc. The study also compares how the summarization attributes affect the results of ROUGE scores. As an extended study, it is revealed that the model of Kuremoto, Tsuruda, Mabu, & Obayashi (2017) is an improved version of the study of Rush et al. (2015) in two-fold. The first one is the proposed model of Kuremoto et al. (2017) carries out the recurrent neural networks instead of the regular neural networks. The second one is that the output text length is adjustable. Then, a sequence-to-sequence model with deep recurrent generative decoder is proposed by Li, Lam, Lidong, & Wang (2017) in which they leveraged the latent structure modeling. The latent variables and the deterministic states play role in the output summaries. Furthermore, a proposal of a coarse-to-fine attention model which improves the speed for long inputs such as documents is presented by Ling & Rush (2017). The idea of this study is to process the whole input with coarse attention to determine the most used parts of the input and to apply fine attention to them while generating the summary.

Another different study is that Nema, Khapra, Laha, & Ravindran (2017) presented a query-based summarization model with two addition models which are query attention model and diversity-based attention model, respectively. Query attention model learns how and when pay attention to text portions. Diversity based attention model tries to alter repetitive portions in output. Moreover, a multi-task learning model that shares decoder parameters with entailment generation model introduced to prevent non-related information in the summary by Pasunuru, Guo, & Bansal (2017). Later, Paulus, Xiong, & Socher (2017) generated a new intra-attentional neural network model that pay attention

over the input text and summarizes separately. Furthermore, the model combines the supervised word prediction with the reinforcement learning. Then, See, Liu, & Manning (2017) proposed regular sequence-to-sequence neural networks with two main additions. Firstly, hybrid pointer generator network was built for copies words from the source text increase the accuracy of extracted information. Secondly, the coverage methodology was proposed to overcome the repetition in the summarized texts. Thus, they proposed the pointer generator sequence-to-sequence neural networks.

Tan, Wan, & Xiao (2017) analyzed the challenging parts of the abstractive text summarization with neural networks. Consequently, graph-based attention model has been proposed in the encoder-decoder model. The outputs of the study demonstrate that the model has considerable results with respect to novel summarization methods. After that, a proposal of sequence-to-sequence model with selective encoding was provided by Zhou, Yang, Wei, & Zhou (2017) where the selective encoding model consists of the sentence encoder, the selective gate networks and the attention-based decoder.

As of 2018, several research on abstractive text summarization has been performed, as well. André, Claudiu, Andreea, & Michael (2018) introduced a diverse beam search approach to improve the neural network loss on the training process. Furthermore, the study also introduced a new metric that measures the extracted text ratio from the input data. Celikyilmaz, Bosselut, He, & Choi (2018) presented an encoder-decoder architecture-based model to summarize the long texts. In the encoder part of the study, the long text is divided into the agents that communicates with each other. Each agent is responsible for a part of a text. Chen & Bansal (2018) proposed two stepped summarization methods named as selecting important sentences and rewriting selected sentences. Those steps are distinct, so to connect steps respectively, sentence level policy gradient method proposed. Furthermore, long text processing speed boosted up thanks to *parallel decoding*. Additionally, Cohan et al. (2018) implemented an early abstractive text summarization model for long inputs such as scientific papers. The proposed model contains novel hierarchical encoder and attentive discourse-aware decoder to construct summary. Later, to explore solutions for poor content selection during the summarization

process a model has generated by Gehrmann, Deng, & Rush (2018). In their study, a data-efficient content selector was carried out to decide phrases that mostly used in the input text. Furthermore, the content selector part combined with the bottom-up attention model led to improvements on the ROUGE results.

Hsu et al. (2018) generated an extractive and abstractive combined summarization model where the sentence-level attention is used to generate word-level attention to eliminate less attended sentence to be used on the summarized text. Furthermore, new inconsistency loss function to achieve more reliable summarized text. Jiang & Bansal (2018) tried to improve the memorization capabilities without using the attention and the pointer mechanisms. Hence, they produced a ‘closed book’ decoder which enforces the encoder part to be more selective. Kryściński, Paulus, Xiong, & Socher (2018) carried out a model to decompose the decoder into a contextual network to reveal related parts of the original text. Moreover, a new metric which is optimized directly to enforce state of the art phrases generation was proposed. Another important study is the study of Li, Bing, & Lam (2018). In this study, an actor-critic based-approach training framework was implemented to increase quality of summaries. The actor is a regular attention-based sequence to sequence framework. The critic consists of combination of the maximum likelihood estimator with global summary quality estimator. Followingly, Li, Xu, Li, & Gao (2018) generated a generation model that combines extractive and abstractive methods. The key concept is the control of the process of generation to prevent the information loss. To achieve that, Key Information Guide Network (KIGN) which encodes the keywords to key information representation has been created. After that, an information controller which positioned between encoder and decoder part was presented by Lin, Sun, Ma, & Su (2018). The introduced controller aims to enhance the quality of the inferences from the input text.

Another significant research has performed by Liu et al. (2018). Liu et al. (2018) propounded two independent models called Generative Model G, Discriminative Model D which are trained simultaneously. Generative model is an instance of reinforcement learning, which takes the input and generates the abstractive summarization.

Discriminative model tries to differentiate the generated summary from the real summary. The study achieves meaningful ROUGE scores over the other novel models on CNN/Daily Mail dataset. Subsequently, Pasunuru & Bansal (2018) introduced two new reward functions via reinforcement learning to achieve a better summarization rate, considering the following three important features such as maintaining saliency, directed logical entailment, and non-redundancy. Furthermore, the study of Wang et al. (2018) proposed a combined model that assembles the convolutional sequence-to-sequence (ConvS2S) model and the self-critical sequence training (SCST) model, and it aims to improve the coherence, the diversity, and the quality of the summarized text.

Regarding the recent studied on the abstractive text summarization methods, Bae, Kim, Kim, & Lee (2019) generated a model that merges the extractive and the abstractive summarization methodologies. The extractive part of the model selects the significant sentences to be used by the abstractive part. Furthermore, the abstractive model has built with the reinforcement learning to increase the summarization quality. In addition, BERT is also integrated into the model. Further, an unsupervised abstractive text summarization model with denoising autoencoder model was produced by Liu, Chung, & Ren (2019). Furthermore, various self-supervised pre-training methods are integrated into model to improve its performance. As a combination of abstractive and extractive methods, Subramanian, Li, Pilault, & Pal (2019) proposed a model that combines extractive and abstractive text summarization methods for the lengthy sequenced texts. In the model, the extractive step is followed by the abstractive model which generates the summarized text. Thanks to model, the summarized text contains less likely copied text from the original text.

Next, Gunel, Zhu, Zeng, & Huang, (2020) proposed a transformer encoder-decoder model that improves coherency in summaries. Liang, Du, & Li, (2020) is aimed generate a solid summary for the social media text (i.e. Twitter) with selective reinforced encoder-decoder model. Liao, Zhang, Chen, & Zhou (2020) carried though a model which is a combination procedure derived from the transform model to resolve the long text

representation. The presented model can hold more memory on encoder part thanks to the analysis on the previously processed information.

As another recent study, (Moirangthem & Lee, 2020) proposed a temporal hierarchical pointer generator to overcome lengthy sequence of texts. Furthermore, Yang et al., (2020) presented a novel model that generates meaningful information, coherent summaries and grammatically correct. Later, transformer-based encoder-decoder with self-supervised objective model was produced by Zhang, Zhao, Saleh, & Liu, (2020).

A summary table for the studies related to the abstractive text summarization algorithms is given in Appendix B in chronological order.

Apart from the aforementioned studies, there are several research articles performing literature survey on text summarization (Bharti & Babu, 2017; Shi, Keneshloo, Ramakrishnan, & K. Reddy 2021). Bharti & Babu (2017) categorized the studies as simple statistical approach, linguistic approach, machine learning approach and hybrid approach in terms of the type of approach. Then, in terms of the type of domain, they have summarized as radio news, journal articles, technical reports, and newspaper articles, etc. Furthermore, Shi, Keneshloo, Ramakrishnan, & K. Reddy (2021) performed a survey on the article regarding their model approach, metric used, and the datasets.

2.2.3 Text Summarization Studies for Turkish

The first study on Turkish text summarization is the study of Altan (2004). Altan (2004) provided a statistical-based system that contains five different modules. The system analyses the words, the sentences, and the paragraphs received in the input document with the pre-calculated weights. Then, Karakaya & Güvenir (2004) proposed a study for the automatic report generation that contains both categorization and summarization. In the initial step of the algorithm, the given document is classified by the given topics. Then the categorized documents are summarized. After that, the summary adds to the generated report. As another early study on Turkish, Uzundere et al. (2008) weighted the sentences according to various features such as title, keywords, positive/negative words, etc. Then,

the summary of texts is obtained based on the users' summary percentages. At the end, 55% success rate has obtained.

One of the leading studies for text summarization in Turkish language is the study of Kutlu, Cıgır, & Cicekli (2010). They proposed a generic text summarization methodology which integrates an extraction method over sentences, where they focus on some important surface-level document features such as terms frequency, title similarity, key phrases, sentence position and centrality. Since this study is one of the leading studies in the Turkish summarization literature, the requirement of data creation was inevitable. Hence, they have generated 2 different datasets: the first one consists of 120 newspaper articles and the other consists of 100 Turkish journal articles. While comparing the performance of the proposed system, ROUGE-1,2, L and W scores are carried out. At the same year, Ozsoy, Cicekli, & Alpaslan (2010) integrated two new algorithms, called Cross and Topic methods, to the latent semantic analysis. They compared the performance of their algorithms with well-known algorithms for text summarization and conclude that Cross method outperforms than the others in terms of ROUGE-L F-measure score.

Güran, Bayazıt, & Gürbüz (2013) introduced a hybrid summarization method for Turkish to incorporate 5 structural features (length, position, title, frequency, class relevance) and 3 semantic features (relevance to each topic, relevance to overall topic, relevance to other sentences). The aim of this study is to achieve an overall score function with the integration of all features. While achieving the overall score function, the weights of the features are obtained by two algorithms: Analytical Hierarchy Process, where weights are calculated by manually, and Artificial Bee Colony algorithms, where the weights are automatically learned. The proposed algorithms are carried out in a newly introduced Turkish corpus (110 news documents) and the results show that the contribution of integration of features is an advantageous way rather than analyzing the features individually.

Hatipoğlu & Omurca (2015) studied a hybrid method where they incorporate statistical scoring sentences and latent semantic analysis. The same authors developed a mobile Turkish summarization system where the structural features (length, position and title),

latent semantic analysis based features (cross method and meaningful word-set method), and Wikipedia based semantic features (counting keywords) are integrated Hatipoglu & Omurca (2016). The cross method that they incorporated is the method of Ozsoy et al. (2010). Next, Demirci, Karabudak, & Ilgen, (2017) studied on producing a single document summary for multi-document news where they carried out the term frequency (TF) to score the documents. Namely, the sentences which have higher importance are to be selected. The system has 43% success rate. As another extraction based text summarization study, Torun & Inner (2018) has also used the term frequency (TF) methodology. In this study, the news articles are summarized, and then the similarity ratio is calculated between the summarized texts.

Later, Doğan & Kaya (2019) proposed an LSTM model followed by a LSA model to summarize the texts. The LSTM part of the study is used for sentiment analysis which is a classification problem. In the study of study Hark, Uckan, Seyyarer, & Karci (2019), they proposed a model that represents text as graphs. In the study, the input sentences are represented as nodes, so edges are interpreted as structure of the input text. Furthermore, (Karakoc & Yılmaz, 2019) studied the abstraction text summarization with an encoder-decoder model. The study is aimed to generate the headlines for the news articles. The model trained with Kemik Haber Dataset and results evaluated with ROUGE-1 score.

In Turkish, 13 studies for the text summarization methods have been analyzed, categorized as provided in Appendix C. Only one abstractive study (Karakoc & Yılmaz, 2019) is encountered in Turkish text summarization literature. Thus, there is a gap in the literature in abstractive text summarization for Turkish.

2.3 Evaluation of Related Works

It has been revealed that the text summarization research area has gaining attention of the researchers all around the world since there is a quite big demand on the text summarization. As a result, academics have been more likely to deal with text summarizing systems in all languages in recent years. Researchers tended to generate advantageous and user-friendly automatic text summarization methods since people want

to reach the summary of texts quickly and instantly while encountering hard to read and long written texts. As seen in Section 2.2, several studies have been conducted mostly in English. 29 studies for the extractive text summarization and 39 studies for the abstractive text summarization methods have been analyzed, categorized, and investigated during the literature review on the summarization in Sections 2.2.1 and 2.2.2, respectively. In addition, 2 survey papers have been explored. The analyzed abstractive and extractive text summarization studies for English and other languages are summarized by indicating their architectures, used data sets and score functions are given in Appendices A and B, respectively.

On the other hand, less studies have been performed on Turkish language, as illustrated in Appendix C, and explained in Section 2.2.3. In Turkish, 13 studies for the text summarization methods have been listed in Appendix C. According to Appendix C, the ones in Turkish language focuses on latent semantic analysis, namely feature extraction methods, while a few of them focuses on encoder-decoder mechanisms. To the best of our knowledge, an abstractive deep neural network method has been performed for Turkish language in only one study (Karakoc & Yılmaz, 2019). Thus, there is a gap in the literature in abstractive text summarization for Turkish.

CHAPTER THREE

TEXT SUMMARIZATION DATASETS AND METHOD

3.1 Datasets

Labeled data set constitutes a substantial part of the supervised machine learning models. In the text summarization models, the labeled data contains the original text and the summarized text. There are a plenty of well conducted studies for English listed in Table 3.1 and well-known datasets are explained briefly in Section 3.1.1. Although there are capable datasets for other tasks in Turkish language, which are demonstrated in Table 3.2, none of them contains the summarized text of the original text. These datasets are explained in Section 3.1.2.

3.1.1 Datasets in English

There are multiple natural language datasets, which serves for the different purposes. The list of the datasets for English together with their resources are presented in Table 3.1.

Table 3.1 Datasets in English

Dataset	Resource URL
AESLC	https://github.com/ryanzhumich/AESLC
Big Patent	https://evasharma.github.io/bigpatent
BillSum	https://github.com/FiscalNote/BillSum
CNN / Daily Mail Dataset	https://github.com/abisee/cnn-dailymail
Cornell Newsroom	http://lil.nlp.cornell.edu/newsroom/index.html
Covid19Sum	https://www.kaggle.com/allen-institute-for-ai/CORD-19-research-challenge
Document Understanding Conferences (DUC) Dataset	https://duc.nist.gov/data.html
Extreme Summarization (XSum) Dataset	https://github.com/EdinburghNLP/XSum/tree/master/XSum-Dataset
Gigaword	https://github.com/harvardnlp/sent-summary
Multi-News	https://github.com/Alex-Fabbri/Multi-News

Table 3.1 Datasets in English (Continues)

Dataset	Resource URL
Newsroom	https://summari.es/
Opinion Abstracts	http://www.ccs.neu.edu/home/luwang/data.html
Opinosis	http://kavita-ganesan.com/opinosis/
Reddit Tifu	https://github.com/ctr4si/MMN
Samsum	https://arxiv.org/src/1911.12237v2/anc
Scientific Papers	https://github.com/armancohan/long-summarization
Text Analysis Conference (TAC) Dataset	https://tac.nist.gov/
Webis-Snippet-20	https://webis.de/data/webis-snippet-20.html
Webis-TLDR-17	https://zenodo.org/record/1168855
WikiHow	https://github.com/mahnazkoupae/WikiHow-Dataset

- **CNN / Daily Mail Dataset:** CNN / Daily Mail Dataset consists of news articles from CNN and Daily Mail with their summarization. This is the most used for text summarization studies among the datasets (“CNN / Daily mail dataset,” n.d.).

- **Document Understanding Conferences (DUC) Dataset:** Document Understanding Conferences (DUC) is a series of conferences had been run annually between 2001 and 2007. The conferences organized by National Institute of Standards and Technology (NIST). These conferences were the forum for the researchers who studies in automatic text summarization to compare their methods and results on the common dataset. The dataset can be requested through their system (“Document understanding conferences datasets - past data,” n.d.).

- **Text Analysis Conference Dataset:** In 2008, Text Analysis Conference (TAC) has been established from Document Understanding Conferences (DUC) and the Question Answering Track of the Text Retrieval Conference (TREC). TAC is also sponsored by NIST. The dataset has the same format with DUC dataset (“Text analysis conference (TAC) dataset,” n.d.).

- **Webis-Snippet-20:** The Webis Abstractive Snippet Corpus 2020 (Webis-Snippet-20) contains 10 million web page/abstractive snippets pairs (“Webis-Snippet-20 dataset,” n.d.).

- **Webis-TLDR-17:** The Webis TLDR Corpus (2017) consists of text and summary pairs from the Reddit dataset. The dataset contains roughly 4 million data collected between 2006-2016. The dataset is one of the first datasets from social media in English. The dataset is targeted for deep learning models (“Dataset for generating TL;DR,” n.d.).

3.1.2 Datasets in Turkish

There are multiple natural language datasets for Turkish, which serves for the different purposes. The list of the datasets for Turkish together with their resources are presented in Table 3.2.

Table 3.2 Datasets in Turkish

Dataset	Resource
TTC-3600: A new benchmark dataset for Turkish text categorization	https://github.com/denopas/TTC-3600
Turkish Datasets from Kemik – Yıldız Technical University	http://www.kemik.yildiz.edu.tr/
Turkish Sentiment Dataset	http://www.baskent.edu.tr/~msert/research/datasets/SentimentDatasetTR.html
English/Turkish Wikipedia Named-Entity Recognition and Text Categorization Dataset	https://data.mendeley.com/datasets/cdcztymf4k/1
TS Corpus	https://tscorpus.com/
1B Tokens Turkish Corpus and Turkish word vectors and analogical reasoning task pairs	https://github.com/onurgu/linguistic-features-in-turkish-word-representations/releases
METU-Sabancı Turkish treebank	http://tools.nlp.itu.edu.tr/Datasets
Turkish Paraphrase Corpus (TuPC)	https://osf.io/wp83a/
Turkish Product Reviews	https://github.com/fthbrmnby/turkish-text-data

- **TTC-3600: A new benchmark dataset for Turkish text categorization:** The dataset has been used for text classification. The dataset has 600 different news articles in 6 different areas such as economy, sport, culture-art, politics, health and technology (“TTC-3600: A new benchmark dataset for Turkish text categorization,” n.d.).

- **Turkish Datasets from Kemik – Yıldız Technical University:** Kemik is a Natural Language Processing group that studies at Yıldız Technical University. The group has various datasets serve in different areas such as classification, n-gram analysis, document similarity (“Turkish Datasets from Kemik – Yıldız Technical University,” n.d.).
- **Turkish Sentiment Dataset:** The data is collected by crawling Turkish tweets to use in sentiment analysis (“Turkish sentiment dataset,” n.d.).
- **English/Turkish Wikipedia Named-Entity Recognition and Text Categorization Dataset:** The dataset is obtained from Turkish and English Wikipedia pages for the purpose of text categorization and named-entity recognition (“English/Turkish Wikipedia named-entity recognition and text categorization dataset,” n.d.).
- **TS Corpus:** TS Corpus is an independent project that contains multiple corpora for the different objectives (“TS Corpus dataset,” n.d.).
- **1B Tokens Turkish Corpus and Turkish word vectors and analogical reasoning task pairs:** The dataset contains word pairs and word vectors (“1B Tokens Turkish corpus and Turkish word vectors and analogical reasoning task pairs,” n.d.).
- **METU-Sabanci Turkish treebank:** The purpose of the dataset is to serve as the test set of the CoNLL-XI shared task (“METU-Sabanci Turkish treebank dataset,” n.d.).
- **Turkish Paraphrase Corpus (TuPC):** The corpus contains 1270 text and paraphrases pairs that had been checked semantically (“Turkish paraphrase corpus,” n.d.).
- **Turkish Product Reviews:** The dataset contains 235000 product reviews from various vendors (“Turkish product reviews,” n.d.).

3.2 Encoder-Decoder Sequence to Sequence (Seq2Seq) Model

Sequence-to-Sequence (Seq2Seq) model is introduced by Sutskever, Vinyals, & Le, (2014) by Google. The idea behind this model is converting the information from one domain into another domain. For example, Google translate uses Seq2Seq models to translate a sentence to another language. The model is useful when the input and output has a fixed length. However, real world examples do not have fixed inputs and outputs,

thus, encoder-decoder architecture has developed. Mostly, the preferred recurrent units in the Seq2Seq models are Long-Short Term Memory (LSTM) or Gated Recurrent Units (GRU) (Sutskever et al., 2014).

The encoder layer is a Recurrent Neural Network (RNN) that process the input sequence and returns an internal state. The output of the encoder layer does not used by the decoder layer; however, it acts as the initial hidden state of the decoder layer. The input words are represented as x_i where i is the order of the word. Hidden states for the neurons are represented as h_i and calculated by Equation (3.1) of Sutskever et al. (2014).

$$h_t = f(W^{(hh)}h_{t-1} + W^{(hx)}x_t) \quad (3.1)$$

The decoder layer is also an RNN that trained to predict the next characters of the target sequence. The decoder uses the hidden state and the cell state of encoder layer to predict the next word in the target sequence. The outputs are represented as y_i which are a collection of words from the answers. The output at any i^{th} step is computed by Equation (3.2). Hidden states for the decoder neurons are computed by Equation (3.3) of Sutskever et al. (2014).

$$y_t = softmax(W^S h_t) \quad (3.2)$$

$$h_t = f(W^{(hh)}h_{t-1}) \quad (3.3)$$

The figure of general Seq2Seq architecture for text translation is depicted in Figure 3.1.

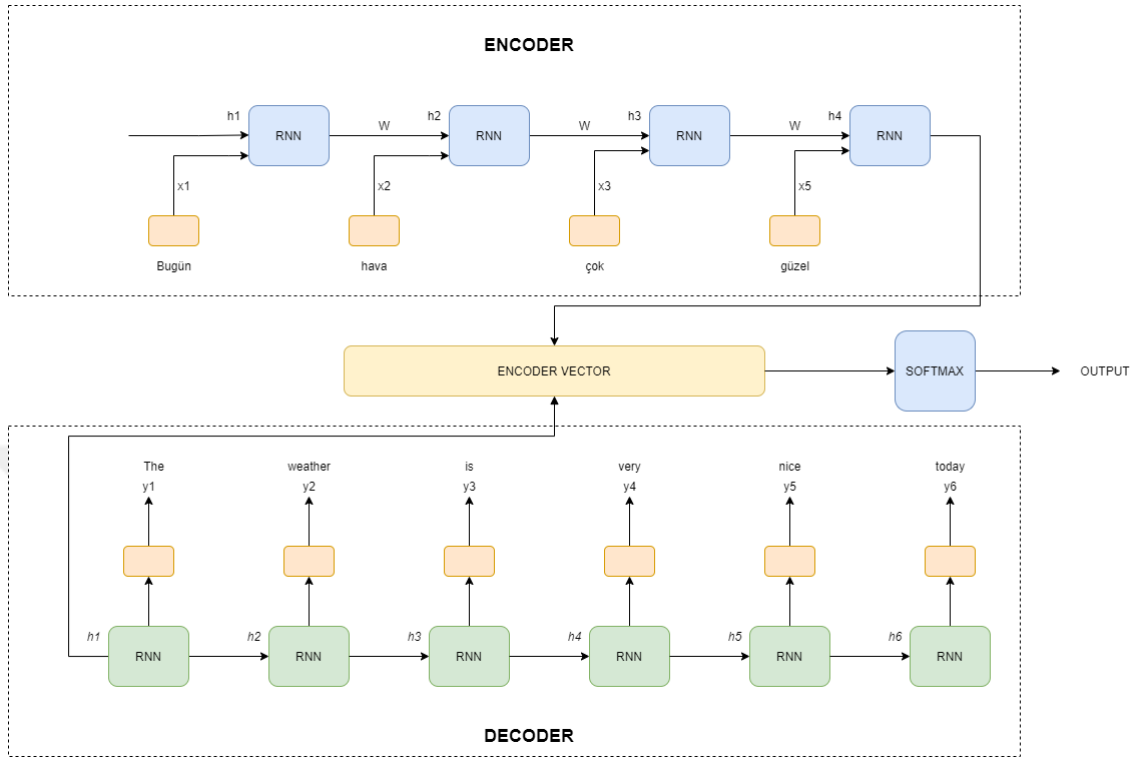


Figure 3.1 General Seq2Seq Architecture for Text Translation

In the figure, the Turkish sentence “Bugün hava çok güzel” is translated to English, as “The weather is very nice today”. As it can be seen from figure, the input and the output sizes are different.

CHAPTER FOUR

TR-SUM: A TEXT SUMMARIZER FOR TURKISH

4.1 General Overview of “TR-Sum: A Text Summarizer for Turkish”

“*TR-Sum: A Text Summarizer for Turkish*” contains three different models that trained and evaluated separately. The general flow of the system is given Figure 4.1.

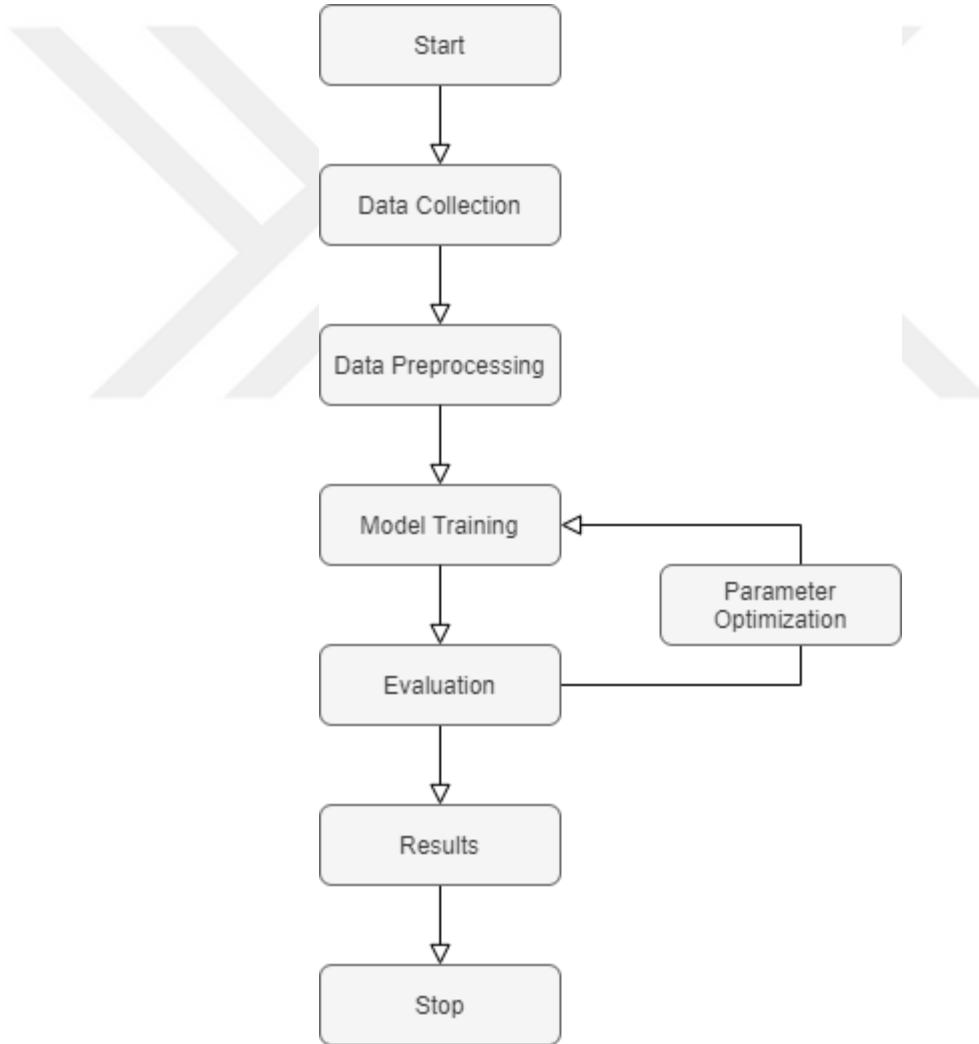


Figure 4.1 Flow Chart of TR-Sum: A Text Summarizer for Turkish

The data collection process and the collected dataset (TR-News-Sum) are described in Section 4.2. Then, the data pre-processing is explained in Section 4.3. Finally, three

different deep neural network models for Turkish are represented in Section 4.4. These models are (i) *Attention Based Seq2Seq Neural Networks*, (ii) *Pointer Generator Seq2Seq Neural Networks* and (iii) *Reinforcement Learning with Seq2Seq Neural Networks*.

4.2 TR-News-Sum Dataset

To be able to implement effective text summarization systems, a dataset constituting original texts and their summaries is a must. In other words, there should be a reference dataset where the summaries for the long-sized texts exist and are reachable so that the text summarization systems can be applied to the dataset. In this thesis, news texts are carried out as the input data. The collaboration had been done with the Journalism Department, Faculty of Communication, Ege University to collect data. A website is developed to collect the news texts and the related summarized news texts from the Journalism Department students. The website is accessible through the link: “<http://txtsum.cs.deu.edu.tr/Account/Login>”. A screenshot of the user interface of the data collection system is given in Figure 4.2.

Metin Özetleme

yigit.yuksel@ceng.deu.edu.trÇıkış Yap

Özetlenecek Metin

Özet Metin

Haber Başlığı

Haber Kategorisi

Siyaset

Haber Linki

Haber Yayınlanma Tarihi

Gönder

Figure 4.2 A Screenshot of the Data Collection System

The data collection system proceeds as follows. The system users (students) are registered by the authorized person, i.e., Professor in the Journalism Department. The users must fill the following mandatory fields: *main text*, *summarized text*, *headline*, *category*, *link (if exists)*, and *publication date* as depicted in Figure 4.2. The collected data saved into the database then exported as csv file format to use for training. Database schema is given in Figure 4.3.

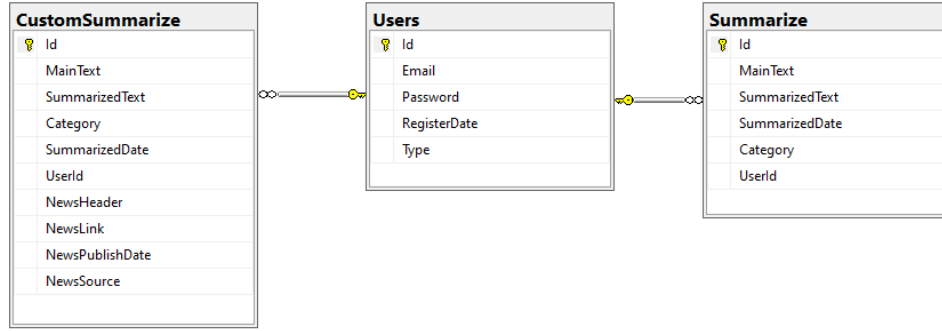


Figure 4.3 Database Schema of the System

150 news texts and their summarizations are collected via the system. These data are used to evaluate and compare the performance of the neural network models proposed in this thesis. In the future, the data collection process can be evaluated as another study to contribute to the Turkish dataset literature.

Five different examples of the collected data from the TR-News-Sum dataset is provided in Appendix D.

4.3 Data Pre-Processing

Preparing data is an important step to get better results from the deep neural network models. Data preprocessing contains two phases: (i) *cleaning*, and (ii) *generating look up dictionary with pre-trained word embeddings*. General flow of the data pre-processing is given Figure 4.4.

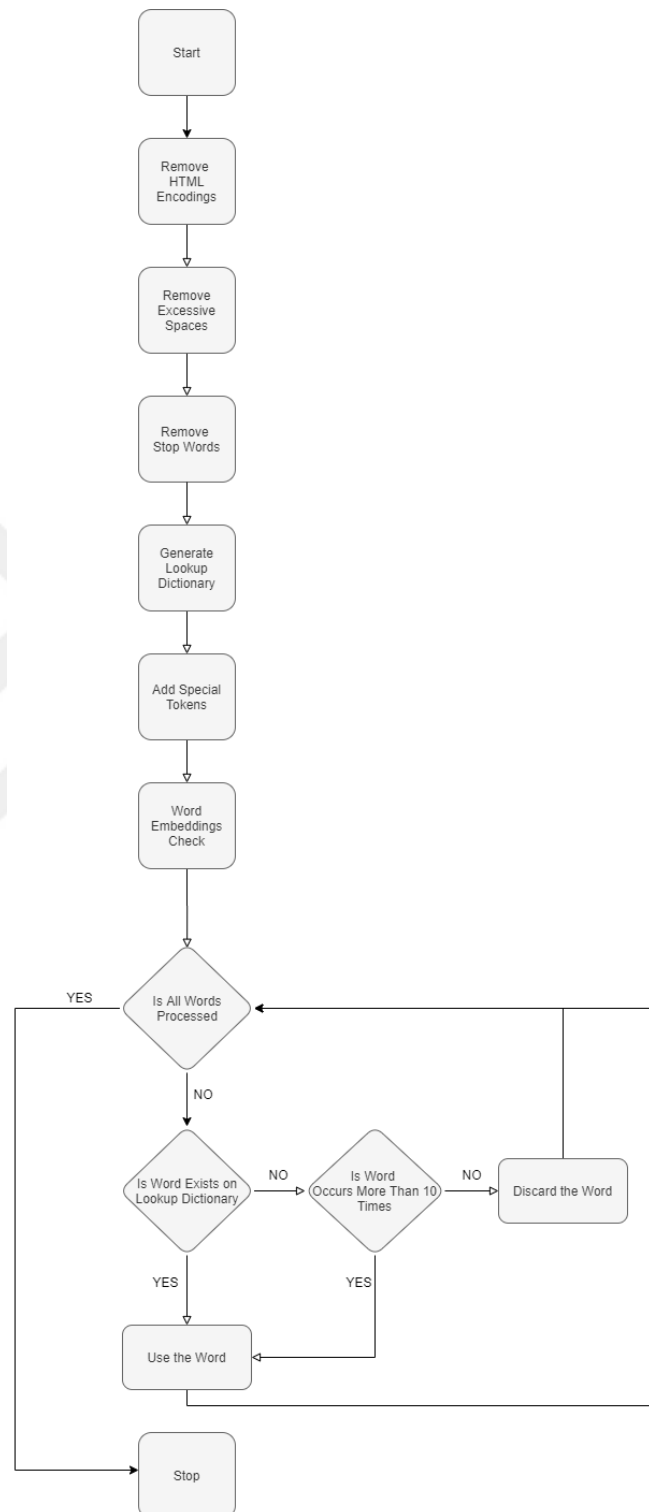


Figure 4.4 Flow Chart of Data Pre-Processing

(i) Cleaning:

This phase is carried out to remove the undesired characters, i.e., HTML encoded characters (i.e., &, %20, etc.), excessive spaces, and stop words. Different stop words are listed in Table 4.1. Namely, it is the formatting phase of the texts. One advantage of this process is that fewer null embeddings will be created.

Table 4.1 Stop Words

{'niye', 'her', 'ama', 'biri', 'ile', 'kez', 'bu', 'niçin', 'yani', 'ya', 'en', 'hiç', 'çok', 'siz', 'nasıl', 'şey', 've', 'birkaç', 'ise', 'aslında', 'da', 'bazı', 'hep', 'acaba', 'nerde', 'de', 'belki', 'biz', 'diye', '3'kim', 'nereye', 'mü', 'o', 'için', 'sanki', 'ki', 'veya', 'mu', 'neden', 'birşey', 'mı', 'daha', 'hem', 'tüm', 'nerede', 'defa', 'çünkü', 'şu', 'hepsi', 'az', 'gibi', 'eğer', 'ne'}

(ii) Generating Look up Dictionary with Pre-trained Word Embeddings:

The neural network models cannot be fed by actual words; thus, words must be converted to the numbers. Initially, the lookup dictionary is generated where each word is represented as a number. The numbers are dependent on the frequency of the word in the corpus. Furthermore, special tokens are added. <EOS> (End of Sentence) and <SOS> (Start of Sentence) tokens are crucial for the neural network models. Additionally, Pad(<PAD>) token is also introduced, because all summaries and texts in a batch must be in same length. Next, the pre-trained word embeddings from *ConceptNet-Numberbatch* and *fastText* are used to increase the training speed and the accuracy.

- **ConceptNet-Numberbatch:** *ConceptNet-Numberbatch* is the part of the *ConceptNet* open data project. *ConceptNet* is a graph that uses the labelled edges to connect sentences and words. *ConceptNet-Numberbatch* is a snapshot of just the word embeddings. *ConceptNet-Numberbatch* contains 51308 unique words for Turkish language. (Speer, Chin, & Havasi, 2017)

- **fastText:** *fastText* is a library for efficient learning of word representations and sentence classification. Pre-trained word embeddings have been generated by *Common Crawl* and *Wikipedia* using *fastText* for 157 languages along with Turkish language. (Mikolov, Grave, Bojanowski, Puhersch, & Joulin, 2017)

Lastly, the input words are constrained by the following criteria: either word exists in the word-embeddings, or the word occurs in the input more than 10 times. The words those do not meet the criteria are tokenized as <UNK> (Unknown Words). The reason for this approach is that less frequent words must be eliminated to increase the summarization quality. The characteristics of the “TR-News-Sum” dataset can be seen Table 4.2. An example for the input can be seen at Table 4.3. The left part of the Table 4.3 shows the input text while the right part of it represents the input vector of the text.

Table 4.2 The Characteristics of the “TR-News-Sum” Dataset

Number of Words	39709
Number of Unique Words	13887
Number of Words Used in Model	3432
Percentage of Used Words	%24.70
Total Number of <UNK>	18341
<UNK> Words Ratio	%46.24

Table 4.3 An Example from Dataset as Input Vector

['<SOS>', 'twitter', 'çalışanlarına', 'koronavirüsün', 'yayılmasını', 'durdurmak', 'için', 'evden', 'çalışmalarını', 'söyledi', 'twitter', 'ın', 'kurucusu', 'jack', 'dorsey', 'uzun', 'zamandır', 'uzaktan', 'çalışma', 'iş', 'modelini', 'desteklediğini', 've', 'kasım', 'ayında', 'bu', 'yılın', 'altı', 'aya', 'kadar', 'olan', 'kısımını', 'afrika', 'da', 'geçirmeyi', 'planladığını', 'açıkladı', '<EOS>']	[397, 208, 2097, 3000, 110, 3001, 5143, 1599, 102, 5144, 8, 153, 186, 312, 5145, 3014, 7, 5146, 287, 2098, 5, 703, 10, 5147, 808, 704, 1601, 254, 1277, 5148, 64, 154, 36, 5149, 2099, 1602, 254, 5150, 5151, 3015, 33, 125, 5152, 2100, 208, 1283, 4, 3016, 2087, 5153, 9, 701, 134, 193, 207, 803, 155, 564, 31, 267, 5154, 5155]
--	---

4.4 The Proposed Neural Network Models for Turkish Text Summarization

Three abstractive text summarization models are proposed in this thesis.

(i) *Attention Based Seq2Seq Neural Networks*

(ii) *Pointer Generator Seq2seq Neural Networks*

(iii) *Reinforcement Learning with Seq2Seq Neural Networks*

(i) Attention Based Seq2Seq Neural Network

Firstly, the idea of Attention Based Seq2Seq Neural Network of Yu, Yue, & Wang (2017) is adapted for Turkish language and applied to the TR-News-Sum dataset. The architecture of the Attention Based Seq2Seq Neural Network is provided in Figure 4.5.

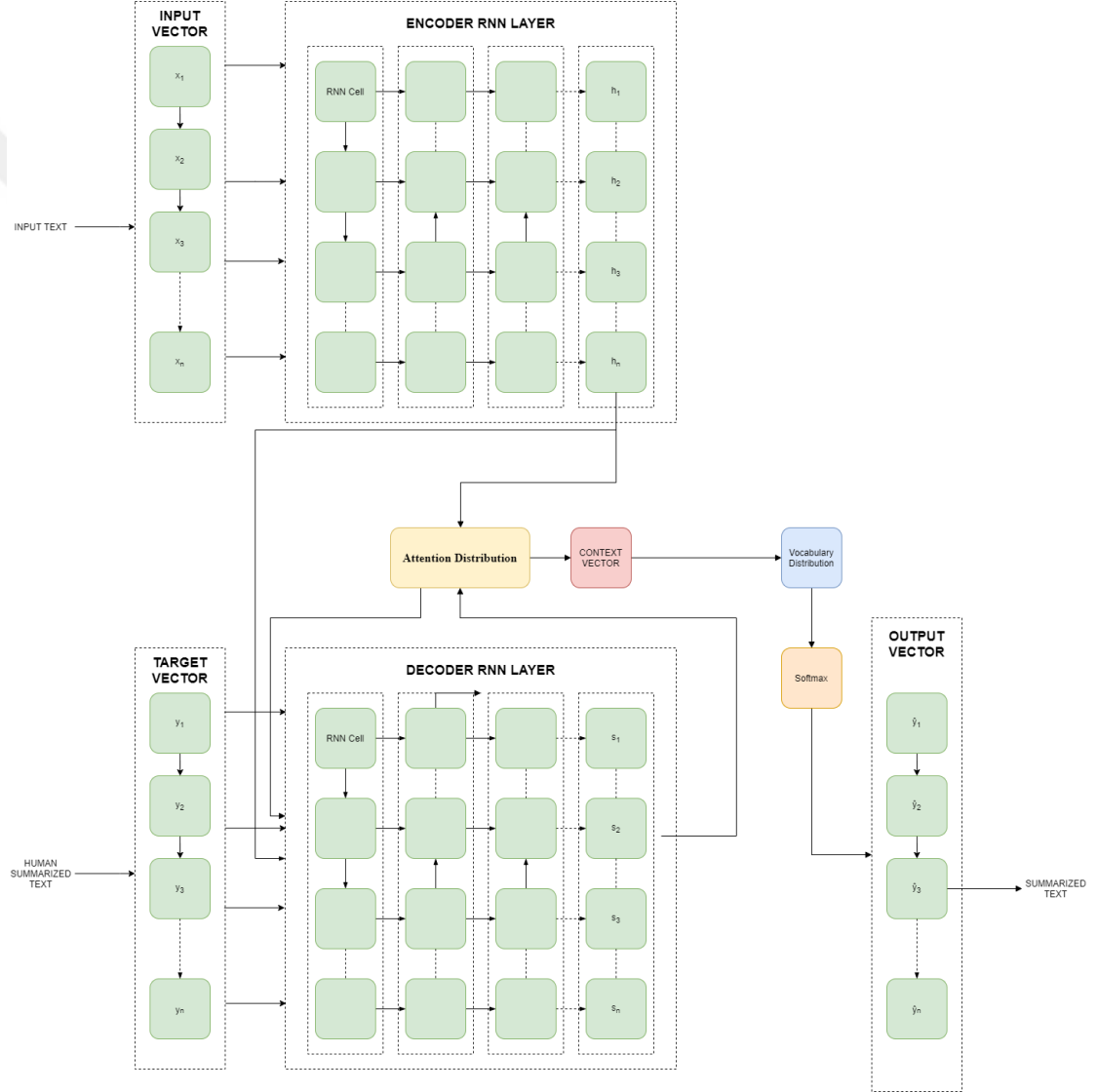


Figure 4.5 Architecture of the Attention Based Seq2Seq Neural Network

In this model, the input data is passed to the encoder layer initially. The output of encoder layer is called as “context vector”. Then, the “context vector” passed on the

decoder layer to generate output sequence. In this step, Bahdanau Attention Mechanism that is provided by Bahdanau, Cho, & Bengio (2014). Instead of attempting to learn a single vector representation for each sentence, the model is trained to focus on various input vectors in the input sequence based on attention weights. During the training process, the decoder layer updates the attention mechanism to adjust weights. These attention weights provide contextual information to the decoder for translation.

(ii) Pointer Generator Seq2Seq Neural Network

As second model, the idea of Pointer Generator Seq2Seq Neural Network model of See et al. (2017) is adapted for Turkish language and applied to the TR-News-Sum dataset. The architecture of the Pointer-Generator Seq2Seq Neural Network is provided in Figure 4.6.

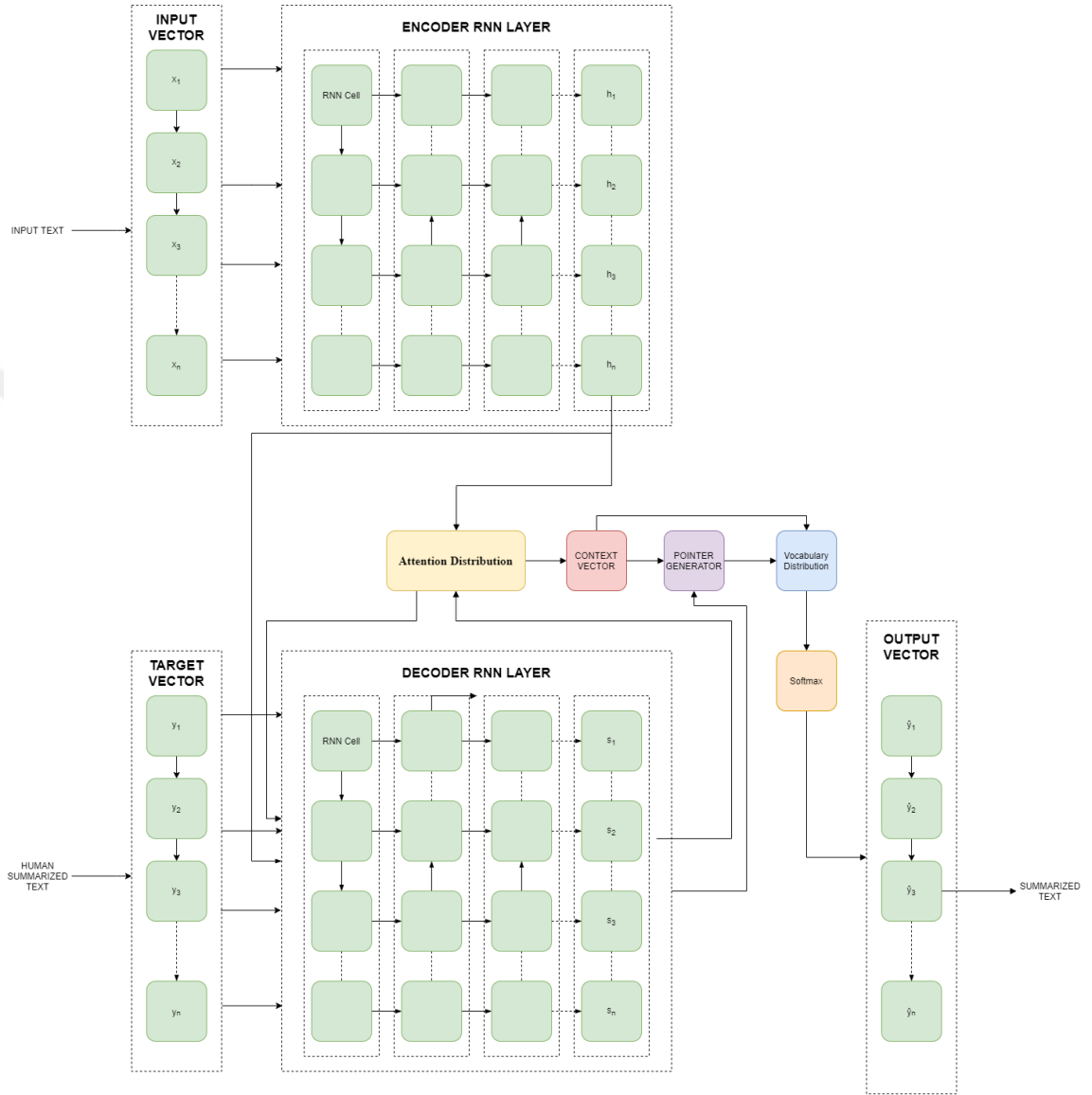


Figure 4.6 Architecture of the Pointer Generator Seq2Seq Neural Network

The model has the option of copying words from the source via pointing while yet being able to generate words from a preset vocabulary. This brings two major benefits to the model. Firstly, the model can handle unseen words while also allows to use a smaller vocabulary. Secondly, the model makes less cost to clone words from the input text. At each decoder step, a generation probability ($P_{gen} \in [0,1]$) is calculated. The decoder

weights are adjusted according to the P_{gen} value. The final distribution, which is the prediction, is weighted and totaled from the vocabulary and attention distributions.

(iii) Reinforcement Learning Based Seq2Seq Neural Network Model (Deep Reinforcement Learning for Sequence-to-Sequence Models)

As another approach, the idea of Reinforcement Learning Based Seq2seq model of Keneshloo, Shi, et al. (2019) is adapted for Turkish language and applied to the TR-News-Sum dataset. The architecture of the Reinforcement Learning Based Seq2Seq Neural Network Model is provided in Figure 4.7.

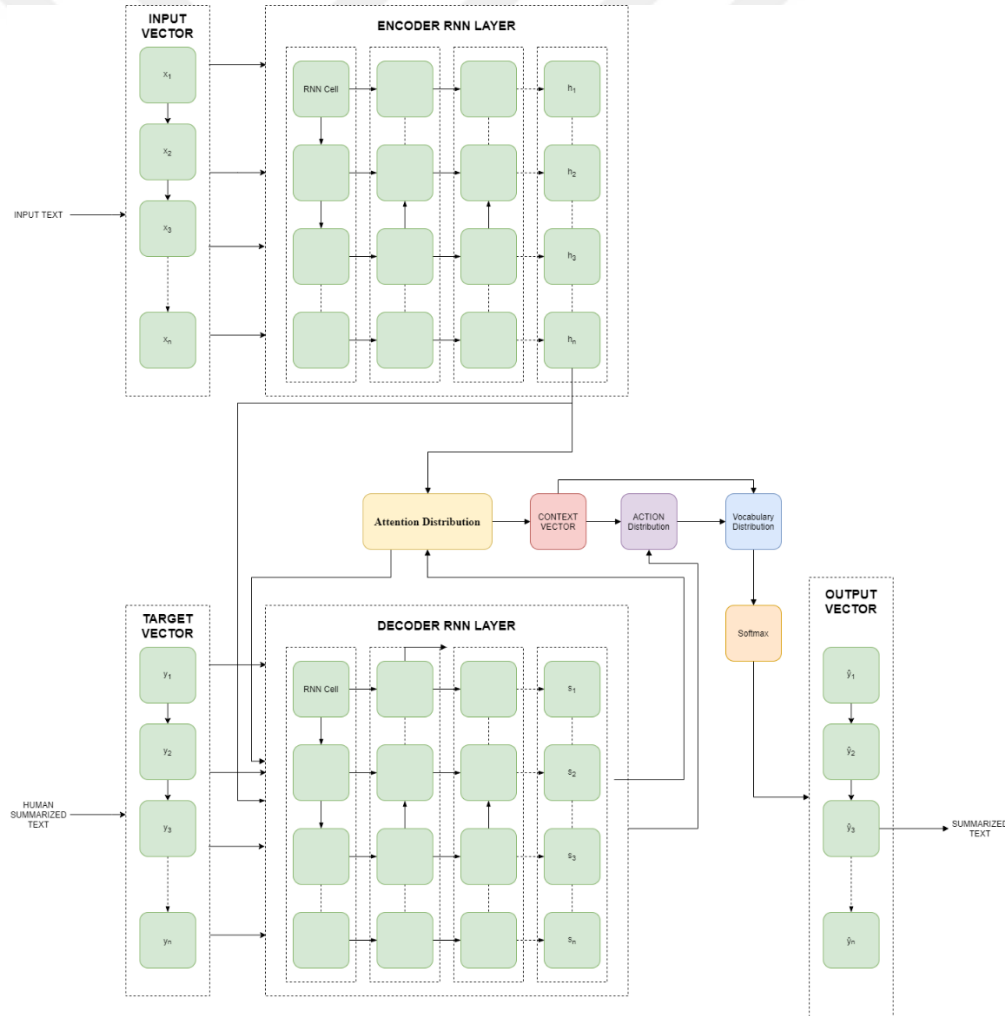


Figure 4.7 Architecture of the Reinforcement Based Seq2Seq Neural Network

In this approach, decision making process and LSTM part of the model is powered by reinforcement learning. The model tries to solve main two problems that occurred on previous models, exposure bias and inconsistency between the train/test measurement, Keneshloo, Shi, et al. (2019). The actor-critic model is designed for the resolve the issues. The actor is a pointer-generator model, and the critic model is a regression model that optimizes action distribution.



CHAPTER FIVE

APPLICATION AND EVALUATION OF TR-SUM: A TEXT SUMMARIZER FOR TURKISH

5.1 Application

The proposed three abstractive text summarization algorithms are applied to the TR-News-Sum dataset, and it is analyzed for Turkish Language. The pseudocode for the application of three models is presented below.

Input: Input news text (X) and user summarized news text (Y).
Output: Trained **Attention Based Seq2Seq Model / Pointer Generator Seq2Seq Neural Network / Reinforcement Learning Based Seq2Seq Neural Network Model**.

Training Steps:

- for** k fold divided input and output sequences (k = 10)
- for** batch of input and output sequences X and Y do
- Run encoding on X and get the last encoder state the
- Run decoding by feeding decoder and obtain the sampled output sequence $\hat{Y}$.
- Calculate the loss according to cross-entropy loss function update the parameters of the model.
- end for**
- end for**

Testing Steps:

- for** k fold divided input and output sequences (k = 10)
- for** batch of input and output sequences X and Y do
- Use the trained model and *argmax* function to sample the output $\hat{Y}$
- Evaluate the model using a performance measure, e.g., ROUGE
- end for**
- end for**

Figure 5.1 Pseudocode of the Models

It is very substantial to mention that the three deep neural network models are studied with two different word embeddings that ConceptNet-Numberbatch and fastText. The main reason is that to use two different word embeddings, compare and measure the effect over the three different models in the Turkish language.

K-Fold cross-validation is used to select the train data and the test data. Moreover, k-fold cross-validation is also helped to avoid overfitting the trained model. K number is selected as 10 which is a common value in machine learning applications. The models are trained on *Google Collaboratory*. The features of the system are as follows.

GPU: 1xTesla K80, compute 3.7, having 2496 CUDA cores, 12GB GDDR5 VRAM
CPU: 1xsingle core hyper threaded Xeon Processors @2.3Ghz i.e. (1 core, 2 threads)
RAM: ~12.6 GB Available
Disk: ~33 GB Available

The application is developed in Python 3.6.9 in Jupyter Notebook. The used libraries are Numpy, Pandas, Tensorflow (version 1.10.1). Furthermore, *Google Colaboratory* has time limitations, thus the model states must be saved on cloud provider, Google Drive. The saved states are used as a knowledge transfer mechanism between the models to increase the summarization quality. The parameters of the application are given in Table 5.1. As it can be seen in the table, the parameters are used for controlling the flow of the system and the performance of the models. The models are trained during 4000 epochs. However, if there are no loss improvement 40 consecutive epochs, the training process is stopped early.

Table 5.1 Parameter Settings

Parameter	Value
Learning rate decay	0.95
Min learning rate	0.0005
Display step	20
Stop early	0
Stop count	4000
Epoch count	3

5.2 Evaluation

To compare the performance of the three deep neural network models, Recall-Oriented Understudy for Gisting Evaluation (ROUGE) scores, which are proposed to the literature by Lin (2004), are carried out. The related studies from the literature, which implements ROUGE scores as their performance metrics, are given in the Appendices A, B and C. The evaluation scores of these studies are presented in the last of column of the Appendices. As it can be seen in the literature review, ROUGE scores are the most

common used performance metrics in the text summarization literature. Hence, in this thesis, the proposed three models are evaluated on the collected Turkish dataset based on the ROUGE-1, ROUGE-2, and ROUGE-L scores. ROUGE-N score computes the *n-gram* co-occurrence statistics between the generated summaries of the models and the reference summaries provided by the dataset. On the other hand, ROUGE-L score computes the longest common sequence between the generated summaries of the model and the reference summaries provided by the dataset. To calculate ROUGE-N and ROUGE-L scores, *recall* and *precision* metrics are required to be calculated. *Recall* metric is calculated by the ratio of the overlapping *n-grams* encountered in both summaries divided by the total number of *n-grams* in the *reference* summary of the dataset. Furthermore, *precision* metric is calculated by the ratio of the overlapping *n-grams* encountered in both summaries divided by the total number of *n-grams* in the *generated* summary of the model. The equations for *recall* and *precision* scores are presented by Equations (5.1) and (5.2), respectively.

$$\begin{aligned} & \frac{\# \text{ of } n - \text{grams found in model and reference}}{\# \text{ of } n - \text{grams in reference}} \\ &= \frac{\text{count}_{\text{match}}(\text{gram}_n)}{\text{count}_{\text{ref}}(\text{gram}_n)} \end{aligned} \quad (5.1)$$

$$\begin{aligned} & \frac{\# \text{ of } n - \text{grams found in model and reference}}{\# \text{ of } n - \text{grams in model}} \\ &= \frac{\text{count}_{\text{match}}(\text{gram}_n)}{\text{count}_{\text{model}}(\text{gram}_n)} \end{aligned} \quad (5.2)$$

Then, ROUGE F1 score can be calculated by Equation (5.3).

$$2 * \frac{\text{precision} * \text{recall}}{\text{precision} + \text{recall}} \quad (5.3)$$

If the precision and recall metrics are computed based on *1-gram*, then the ROUGE-1 score is obtained. Similarly, if the precision and recall metrics are computed based on *2-grams* then the ROUGE-2 score is obtained.

To calculate ROUGE-L score, recall and precision metrics must be updated as follows. Instead of calculating the number of *n-grams* found both in model and reference in the nominator, the *longest common sequence* found in model and reference is calculated. The *longest common sequence* is the common subsequence of two given sequences with maximum length (Lin, 2004). Hence, the recall and precision metrics are updated by Equations (5.4) and (5.5).

$$\frac{\text{Longest common subsequence}}{\# \text{ of } n - \text{grams in reference}} = \frac{LCS(gram_n)}{count_{ref}(gram_n)} \quad (5.4)$$

$$\frac{\text{Longest common subsequence}}{\# \text{ of } n - \text{grams in model}} = \frac{LCS(gram_n)}{count_{model}(gram_n)} \quad (5.5)$$

Then, the ROUGE-L score can be calculated by Equation (5.3) with recall and precision values obtained by Equations (5.4) and (5.5).

The average ROUGE-1, ROUGE-2 and ROUGE-L scores for each model by both word embeddings is reported in Table 5.2.

Table 5.2 Average Rouge Scores of the Three Neural Network Models

Models Avg Rouge Scores	<i>(i)</i> Attention Based Seq2Seq Neural Network		<i>(ii)</i> Pointer Generator Seq2Seq Neural Network		<i>(iii)</i> Reinforcement Learning Based Seq2Seq Neural Network	
Word Embedding Source	ConceptNet-Numberbatch	fastText	ConceptNet-Numberbach	fastText	ConceptNet-Numberbatch	fastText
Rouge-1	0.146	0.229	0.147	0.304	0.136	0.283
Rouge-2	0.097	0.281	0.102	0.352	0.090	0.334
Rouge-L	0.112	0.443	0.122	0.477	0.117	0.464

Three different examples of the three proposed deep neural network models from the TR-News-Sum dataset are provided in Appendix E. Three examples are the good performing examples based on Pointer Generator Seq2Seq Neural Network. In Appendix E, firstly, the collected summarized texts are provided. Then, the summarized text obtained by the three neural network models which are studied with both *ConceptNet-Numberbatch* and *fastText* word embeddings are provided for three different examples.

As it can be seen in Table 5.2, with *ConceptNet-Numberbatch* word embedding, Attention Based Seq2Seq Neural Network has values of 0.146, 0.097, and 0.112 for ROUGE-1, ROUGE-2, and ROUGE-L scores, respectively. Pointer Generator Seq2Seq Neural Network has values of 0.147, 0.102, and 0.122 for ROUGE-1, ROUGE-2, and ROUGE-L scores, respectively. Lastly, Reinforcement Learning Based Seq2Seq Neural Network has values of 0.136, 0.090, and 0.117 for ROUGE-1, ROUGE-2 and ROUGE-L scores, respectively. Although ROUGE-1 scores are the highest ones among all ROUGE scores for each model that is studied with *ConceptNet-Numberbatch* word embedding, ROUGE-2 and ROUGE-L scores are quite low. Thus, there is enough room to improve

the performance of the models when are studied with *ConceptNet-Numberbatch* word embedding.

As it can be seen in Table 5.2, with *fastText* word embedding, Attention Based Seq2Seq Neural Network has values of 0.229, 0.281, and 0.443 for ROUGE-1, ROUGE-2, and ROUGE-L scores, respectively. Pointer Generator Seq2Seq Neural Network has values of 0.304, 0.352, and 0.477 for ROUGE-1, ROUGE-2, and ROUGE-L scores, respectively. Lastly, Reinforcement Learning Based Seq2Seq Neural Network has values of 0.283, 0.334, and 0.464 for ROUGE-1, ROUGE-2 and ROUGE-L scores, respectively. All ROUGE scores are of good quality for each model that is studied with *fastText* word embedding. Furthermore, ROUGE-L scores are the highest ones among all ROUGE scores.

It can be inferred from Table 5.2 that each ROUGE score of models that are achieved by *ConceptNet-Numberbatch* word embedding is improved if the models are studied with *fastText* word embedding. The improvement on all ROUGE scores achieved by *fastText* word embedding for each neural network model is depicted in Table 5.3. Namely, the differences between the columns for each neural network model are reported.

Table 5.3 Improvement Achieved by *fastText* Word Embedding on Three Neural Network Models

Models Avg Rouge Scores	(i) Attention Based Seq2Seq Neural Network	(ii) Pointer Generator Seq2Seq Neural Network	(iii) Reinforcement Learning Based Seq2Seq Neural Network
Rouge-1	0.083	0.157	0.147
Rouge-2	0.184	0.250	0.244
Rouge-L	0.331	0.355	0.347

The improvements on all ROUGE scores achieved by *fastText* word embedding for each neural network model can be concluded as follows.

- The ROUGE-1, ROUGE-2 and ROUGE-L scores of *Attention Based Seq2Seq Neural Network* model are improved by 0.083, 0.184 and 0.331 units, respectively, with the implementation of *fastText* word embedding.
- The ROUGE-1, ROUGE-2 and ROUGE-L scores of *Pointer Generator Seq2Seq Neural Network* model are improved by 0.157, 0.250 and 0.355 units, respectively, with the implementation of *fastText* word embedding.
- The ROUGE-1, ROUGE-2 and ROUGE-L scores of *Reinforcement Learning Based Seq2Seq Neural Network* model are improved by 0.147, 0.244 and 0.347 units, respectively, with the implementation of *fastText* word embedding.

To conclude, the overall inferences from the computational experimentation are:

- Three deep neural network models that are trained by *ConceptNet-Numberbatch* word embedding produce lower ROUGE scores.
- However, all ROUGE scores of the three deep neural network models trained by *ConceptNet-Numberbatch* word embedding are improved good enough with the inclusion of *fastText* word embedding into the models.
- All ROUGE scores of each model that is studied with *fastText* word embedding have decent quality Also, ROUGE-L scores are the highest ones among all ROUGE scores for each model that is studied with *fastText* word embedding.
- The highest ROUGE scores are acquired by *Pointer Generator Seq2Seq Neural Network* with *fastText* word embedding. Particularly, *Pointer Generator Seq2Seq Neural Network* has a value of 0.477 for ROUGE-L score that is extremely high which leads us to conclude that *Pointer Generator Seq2Seq Neural Network* is a promising deep neural network model that may produce text summaries with good quality.

CHAPTER SIX

CONCLUSION AND FUTURE WORK

According to the research and review performed on the literature so far, an abstractive text summarization system for Turkish is a very significant research direction. A web interface of the data collection system is developed to collect data from users. Then, a news dataset named as “TR-News-Sum” dataset is generated with the collected data. The collected dataset is preprocessed to be feed to the deep neural network models. Following, three deep neural network models are adapted for Turkish language and studied on the “TR-News-Sum” dataset. These models are (i) *Attention Based Seq2Seq Neural Networks*, (ii) *Pointer Generator Seq2Seq Neural Networks* and (iii) *Reinforcement Learning with Seq2Seq Neural Networks*. All three models are preprocessed with both *ConceptNet-Numberbatch* word embeddings and *fastText* word embeddings. K-Fold cross-validation is used to select the train data and the test data. To compare the performance of the three deep neural network models, Recall-Oriented Understudy for Gisting Evaluation (ROUGE) scores, which are proposed to the literature by Lin (2004), are carried out. The proposed three models are evaluated on the collected Turkish dataset based on the ROUGE-1, ROUGE-2, and ROUGE-L scores. Hence, an automatic abstractive text summarization method for Turkish language, named as “*TR-Sum: A Text Summarizer for Turkish*”, is developed. The average ROUGE-1, ROUGE-2 and ROUGE-L scores for each model by both word embeddings is reported in Table 6.1.

Table 6.1 Average Rouge Scores of the Three Neural Network Models

Models Avg Rouge Scores	<i>(i)</i> Attention Based Seq2Seq Neural Network		<i>(ii)</i> Pointer Generator Seq2Seq Neural Network		<i>(iii)</i> Reinforcement Learning Based Seq2Seq Neural Network	
Word Embedding Source	ConceptNet-Numberbatch	fastText	ConceptNet-Numberbach	fastText	ConceptNet-Numberbatch	fastText
Rouge-1	0.146	0.229	0.147	0.304	0.136	0.283
Rouge-2	0.097	0.281	0.102	0.352	0.090	0.334
Rouge-L	0.112	0.443	0.122	0.477	0.117	0.464

The overall inferences that can be reported from Table 6.3 are:

- Three deep neural network models that are trained by *ConceptNet-Numberbatch* word embedding produce lower ROUGE scores.
- However, all ROUGE scores of the three deep neural network models trained by *ConceptNet-Numberbatch* word embedding are improved good enough with the inclusion of *fastText* word embedding into the models.
- All ROUGE scores of each model that is studied with *fastText* word embedding have decent quality. Also, ROUGE-L scores are the highest ones among all ROUGE scores for each model that is studied with *fastText* word embedding.
- The highest ROUGE scores are acquired by *Pointer Generator Seq2Seq Neural Network* with *fastText* word embedding. Particularly, *Pointer Generator Seq2Seq Neural Network* has a value of 0.477 for ROUGE-L score that is extremely high which leads us to conclude that *Pointer Generator Seq2Seq Neural Network* is a promising deep neural network model that may produce text summaries with good quality.

During the progress of this thesis, one of the challenging parts of the study was the data collection. It is hard to generate enough number of summaries of the main texts. To

overcome this issue, a collaboration had been performed with the Journalism Department, Faculty of Communication, Ege University to collect data. Another problem was to find a suitable development environment to train the models. Because the training process is GPU intensive operation. To overcome this issue, cloud computing, *Google Collaboratory*, is used in the training process of the models.

The contributions of this thesis can be listed as follows.

- i. The related works in the literature of extractive and abstractive text summarization are analyzed. Then, the works for the Turkish language are also explored.
- ii. A web interface of the data collection system is constituted to collect data from users.
- iii. A novel news dataset named as “TR-News-Sum” dataset is generated with the collected data.
- iv. The collected dataset is preprocessed to be feed to the deep neural network models. Two different word embeddings, *ConceptNet-Numberbatch* and *fastText* word embeddings, are used during the data pre-processing.
- v. Following, three deep neural network models are studied on the “TR-News-Sum” dataset. These models are (i) *Attention Based Seq2Seq Neural Networks*, (ii) *Pointer Generator Seq2Seq Neural Networks* and (iii) *Reinforcement Learning with Seq2Seq Neural Networks*. All three models are preprocessed with both *ConceptNet-Numberbatch* word embeddings and *fastText* word embeddings.
- vi. These models are evaluated on the collected Turkish dataset based on the ROUGE-1, ROUGE-2, and ROUGE-L scores. According to the computational experimentation, *Pointer Generator Seq2Seq Neural Network* is a promising deep neural network model that may produce text summaries with good quality for Turkish language.

There are possible future research directions of this thesis. One of them is to improve the generated dataset with the addition of new texts and their summaries. More summaries

are required to increase the performance of the models. In addition, as it can be seen from the dataset, the existed summaries generally consist of the first two sentences of the main texts. Hence, increasing the quality of these summaries will consecutively improve the performance of the models. The more unique summaries bring about the more improved models. Moreover, different models can be also studied to generate more meaningful summaries in Turkish language. Not but not least, the existed words embeddings can be adapted for the state-of-art neural network models or new word embeddings can be generated.



REFERENCES

- 1B Tokens Turkish corpus and Turkish word vectors and analogical reasoning task pairs.* (n.d.). Retrieved August 9, 2021, from <https://github.com/onurgu/linguistic-features-in-turkish-word-representations/releases>
- Al-Sabahi, K., Zuping, Z., & Nadher, M. (2018). A hierarchical structured self-attentive model for extractive document summarization (HSSAS). *IEEE Access*, 6, 24205–24212.
- Altan, Z. (2004). A Turkish automatic text summarization system. In *Proceedings of the IASTED Internatoinal Conference Artificial Intelligence and Applications, February 16-18*.
- Amancio, D. R., Nunes, M. G. V., Oliveira, O. N., & Costa, L. D. F. (2012). Extractive summarization using complex networks and syntactic dependency. *Physica A: Statistical Mechanics and Its Applications*, 391(4), 1855–1864.
- André, C., Claudiu, M., Andreea, H., & Michael, B. (2018). Diverse beam search for increased novelty in abstractive summarization. *ArXiv Preprint ArXiv:1802.01457*.
- Bae, S., Kim, T., Kim, J., & Lee, S. (2019). Summary level training of sentence rewriting for abstractive summarization. *ArXiv Preprint ArXiv:1909.08752*.
- Bahdanau, D., Cho, K., & Bengio, Y. (2014). Neural machine translation by jointly learning to align and translate. *ArXiv Preprint ArXiv:1409.0473*.
- Barzilay, R., & McKeown, K. R. (2005). Sentence fusion for multidocument news summarization. *Computational Linguistics*, 31(3), 297–327.
- Bhargava, R., & Sharma, Y. (2020). Deep extractive text summarization. In *Procedia Computer Science* (Vol. 167, pp. 138–146). Elsevier B.V.
- Bharti, S. K., & Babu, K. S. (2017). Automatic keyword extraction for text summarization: A survey. *ArXiv Preprint ArXiv:1704.03242*.
- Brito, E., Lübbering, M., Biesner, D., Hillebrand, L. P., & Bauckhage, C. (2019). Towards

- supervised extractive text summarization via RNN-based sequence classification. *ArXiv Preprint ArXiv:1911.06121*.
- Celikyilmaz, A., Bosselut, A., He, X., & Choi, Y. (2018). Deep communicating agents for abstractive summarization. *NAACL HLT 2018 - 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies - Proceedings of the Conference*, 1, 1662–1675.
- Chen, Q., Zhu, X., Ling, Z., Wei, S., & Jiang, H. (2016). Distraction-based neural networks for document summarization. *ArXiv Preprint ArXiv:1610.08462*.
- Chen, X., Gao, S., Tao, C., Song, Y., Zhao, D., & Yan, R. (2018). Iterative document representation learning towards summarization with polishing. *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing, EMNLP 2018*, 4088–4097.
- Chen, Y.-C., & Bansal, M. (2018). Fast abstractive summarization with reinforce-selected sentence rewriting. *ACL 2018 - 56th Annual Meeting of the Association for Computational Linguistics, Proceedings of the Conference (Long Papers)*, 1, 675–686.
- CNN / Daily mail dataset*. (n.d.). Retrieved August 9, 2021, from <https://github.com/abisee/cnn-dailymail>
- Cohan, A., Deroncourt, F., Kim, D. S., Bui, T., Kim, S., Chang, W., & Goharian, N. (2018). A discourse-aware attention model for abstractive summarization of long documents. *NAACL HLT 2018 - 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies - Proceedings of the Conference*, 2, 615–621.
- Dataset for generating TL;DR*. (n.d.). Retrieved August 9, 2021, from <https://zenodo.org/record/1168855>
- Demirci, F., Karabudak, E., & Ilgen, B. (2017). Multi-document summarization for Turkish news. In *2017 International Artificial Intelligence and Data Processing*

Symposium 1–5.

Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K. (2019). BERT: Pre-training of deep bidirectional transformers for language understanding. *Proceedings of the 2019 Conference of the North*, 4171–4186.

Document understanding conferences datasets - past data. (n.d.). Retrieved August 9, 2021, from <https://duc.nist.gov/data.html>

Doğan, E., & Kaya, B. (2019). Deep learning based sentiment analysis and text summarization in social networks. In *2019 International Conference on Artificial Intelligence and Data Processing Symposium (IDAP), Malatya, Turkey, 2019* 1–6.

Elbarougy, R., Behery, G., & El Khatib, A. (2020). Extractive Arabic text summarization using modified PageRank algorithm. *Egyptian Informatics Journal*, 21(2), 73–81.

English/Turkish Wikipedia named-entity recognition and text categorization dataset. (n.d.). Retrieved August 9, 2021, from <https://data.mendeley.com/datasets/cdcztymf4k/1%0D%0A>

Erkan, G., & Radev, D. R. (2004). LexRank: Graph-based lexical centrality as salience in text summarization. *Journal of Artificial Intelligence Research*, 22, 457–479.

Fan, A., Grangier, D., & Auli, M. (2017). Controllable abstractive summarization. In *arXiv preprint arXiv:1711.05217*.

Fang, C., Mu, D., Deng, Z., & Wu, Z. (2017). Word-sentence co-ranking for automatic extractive text summarization. *Expert Systems with Applications*, 72, 189–195.

Ferreira, R., De Souza Cabral, L., Freitas, F., Lins, R. D., De França Silva, G., Simske, S. J., & Favaro, L. (2014). A multi-document summarization system based on statistics and linguistic treatment. *Expert Systems with Applications*, 41(13), 5780–5787.

Gehrmann, S., Deng, Y., & Rush, A. M. (2018). Bottom-up abstractive summarization. *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing, EMNLP 2018*, 4098–4109.

- Ghodratnama, S., Beheshti, A., Zakershaharak, M., & Sobhanmanesh, F. (2020). Extractive document summarization based on dynamic feature space mapping. *IEEE Access*, 8, 139084–139095.
- Gunel, B., Zhu, C., Zeng, M., & Huang, X. (2020). Mind the facts: Knowledge-noosted coherent abstractive text summarization. *ArXiv Preprint ArXiv:2006.15435*.
- Güran, A., Bayazıt, N. G., & Gürbüz, M. Z. (2013). Efficient feature integration with Wikipedia-based semantic feature extraction for Turkish text summarization. *Turkish Journal of Electrical Engineering & Computer Sciences*, 21(5), 1411–1425.
- Hark, C., Uckan, T., Seyyarer, E., & Karci, A. (2019). Extractive text summarization via graph entropy (Çizge entropi ile çıkarıcı metin özetleme). In *2019 International Conference on Artificial Intelligence and Data Processing Symposium (IDAP), Malatya, Turkey, 2019* 1–5.
- Hatipoğlu, A., & Omurca, S. İ. (2016). A Turkish Wikipedia text summarization system for mobile devices. *Information Technology and Computer Science*, 01, 1–10. <https://doi.org/10.5815/ijitcs.2016.01.01>
- Hatipoğlu, A., & Omurca, S. İ. (2015). Türkçe metin özetlemede melez modelleme (A hybrid modelling for Turkish text summarization). *Dokuz Eylül Üniversitesi Mühendislik Fakültesi Fen ve Mühendislik Dergisi*, 17(50), 95–108.
- Hernandez-Castaneda, A., Garcia-Hernandez, R. A., Ledeneva, Y., & Millan-Hernandez, C. E. (2020). Extractive automatic text summarization based on lexical-semantic keywords. *IEEE Access*, 8, 49896–49907.
- Hsu, W.-T., Lin, C.-K., Lee, M.-Y., Min, K., Tang, J., & Sun, M. (2018). A unified model for extractive and abstractive summarization using inconsistency loss. *ACL 2018 - 56th Annual Meeting of the Association for Computational Linguistics, Proceedings of the Conference (Long Papers)*, 1, 132–141.
- Jadhav, A., & Rajan, V. (2018). Extractive summarization with SWAP-NET: Sentences and words from alternating pointer networks. In *Proceedings of the 56th annual*

meeting of the association for computational linguistics (pp. 142–151).

- Jain, A., Bhatia, D., & Thakur, M. K. (2017). Extractive text summarization using word vector embedding. In *Proceedings - 2017 International Conference on Machine Learning and Data Science, MLDS 2017* (Vol. 2018-Janua, pp. 51–55). Institute of Electrical and Electronics Engineers Inc.
- Jiang, Y., & Bansal, M. (2018). Closed-book training to improve summarization encoder memory. *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing, EMNLP 2018*, 4067–4077.
- Karakaya, K. M., & Güvenir, H. A. (2004). ARG: A tool for automatic report generation. *Istanbul University - Journal of Electrical & Electronics Engineering*, 4.2, 1101–1109.
- Karakoc, E., & Yilmaz, B. (2019). Deep learning based abstractive Turkish news summarization. In *27th Signal Processing and Communications Applications Conference (SIU) , IEEE*, 1–4.
- Keneshloo, Y., Ramakrishnan, N., & Reddy, C. K. (2019). Deep transfer reinforcement learning for text summarization. In *SIAM International Conference on Data Mining, SDM 2019*, 675–683. Society for Industrial and Applied Mathematics Publications.
- Keneshloo, Y., Shi, T., Ramakrishnan, N., & Reddy, C. K. (2019). Deep reinforcement learning for sequence-to-sequence models. *IEEE Transactions on Neural Networks and Learning Systems*, 31(7), 2469–2489.
- Khatri, C., Singh, G., & Parikh, N. (2018). Abstractive and extractive text summarization using document context vector and recurrent neural networks. *ArXiv Preprint ArXiv:1807.08000*.
- Kryściński, W., Paulus, R., Xiong, C., & Socher, R. (2018). Improving abstraction in text summarization. *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing, EMNLP 2018*, 1808–1817.

- Kuremoto, T., Tsuruda, T., Mabu, S., & Obayashi, M. (2017). A sentence summarizer using recurrent neural network and attention-based encoder. In *Proceedings of 2017 International Conference on Applied Mathematics, Modeling and Simulation (AMMS 2017)*, 245–248. Atlantis Press.
- Kutlu, M., Cıgır, C., & Cicekli, I. (2010). Generic text summarization for Turkish. *The Computer Journal*, 53(8), 1315–1323.
- Li, C., Xu, W., Li, S., & Gao, S. (2018). Guiding generation for abstractive text summarization based on key information guide network. In *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, 55–60.
- Li, P., Bing, L., & Lam, W. (2018). Actor-critic based training framework for abstractive summarization. *ArXiv Preprint ArXiv:1803.11070*.
- Li, P., Lam, W., Lidong, B., & Wang, Z. (2017). Deep recurrent generative decoder for abstractive text summarization. In *arXiv preprint arXiv:1708.00625*.
- Liang, Z., Du, J., & Li, C. (2020). Abstractive social media text summarization using selective reinforced Seq2Seq attention model. *Neurocomputing*, 410, 432–440.
- Liao, P., Zhang, C., Chen, X., & Zhou, X. (2020). Improving abstractive text summarization with history aggregation. In *Proceedings of the International Joint Conference on Neural Networks*, 1–9. Institute of Electrical and Electronics Engineers Inc.
- Lin, C.-Y. (2004). ROUGE: A Package for automatic evaluation of summaries. *Text Summarization Branches Out*, 74–81.
- Lin, J., Sun, X., Ma, S., & Su, Q. (2018). Global encoding for abstractive summarization. *ACL 2018 - 56th Annual Meeting of the Association for Computational Linguistics, Proceedings of the Conference (Long Papers)*, 2, 163–169.
- Ling, J., & Rush, A. M. (2017). *Coarse-to-fine attention models for document*

- summarization*. PhD Thesis, Harvard University, Cambridge.
- Liu, L., Lu, Y., Yang, M., Qu, Q., Zhu, J., & Li, H. (2018). Generative adversarial network for abstractive text summarization. In *Proceedings of the Thirty-second AAAI Conference on Artificial Intelligence*, 32(1), 8109–8110.
- Liu, P. J., Chung, Y.-A., & Ren, J. (2019). SummAE: Zero-shot abstractive text summarization using length-agnostic auto-encoders. *ArXiv Preprint ArXiv:1910.00998*.
- Liu, Y., & Lapata, M. (2019). Text summarization with pretrained encoders. *EMNLP-IJCNLP 2019 - 2019 Conference on Empirical Methods in Natural Language Processing and 9th International Joint Conference on Natural Language Processing, Proceedings of the Conference*, 3730–3740.
- Lopyrev, K. (2015). Generating news headlines with recurrent neural networks. *ArXiv Preprint ArXiv:1512.01712*.
- Madhuri, J. N., & Ganesh Kumar, R. (2019). Extractive text summarization using sentence ranking. In *2019 International Conference on Data Science and Communication, IconDSC 2019*. Institute of Electrical and Electronics Engineers Inc.
- METU-Sabanci Turkish treebank dataset*. (n.d.). Retrieved August 9, 2021, from <http://tools.nlp.itu.edu.tr/Datasets>
- Mikolov, T., Grave, E., Bojanowski, P., Puhersch, C., & Joulin, A. (2017). Advances in Pre-Training Distributed Word Representations. Retrieved from <https://commoncrawl.org/2017/06>
- Mohd, M., Jan, R., & Shah, M. (2020). Text document summarization using word embedding. *Expert Systems with Applications*, 143, 112958. <https://doi.org/10.1016/J.ESWA.2019.112958>
- Moirangthem, D. S., & Lee, M. (2020). Abstractive summarization of long texts by representing multiple compositionality with temporal hierarchical pointer generator

network. *Neural Networks*, 124, 1–11.
<https://doi.org/10.1016/J.NEUNET.2019.12.022>

Nallapati, R., Zhai, F., & Zhou, B. (2017). SummaRuNNer: A recurrent neural network based sequence model for extractive summarization of documents. In *Thirty-First AAAI Conference on Artificial Intelligence*, 3075–3081.

Nallapati, R., Zhou, B., Santos, C. dos, & Xiang, B. (2016). Abstractive text summarization using sequence-to-sequence RNNs and beyond. *ArXiv Preprint ArXiv:1602.06023*.

Nema, P., Khapra, M., Laha, A., & Ravindran, B. (2017). Diversity driven attention model for query-based abstractive summarization. *ACL 2017 - 55th Annual Meeting of the Association for Computational Linguistics, Proceedings of the Conference (Long Papers)*, 1, 1063–1072.

Ozsoy, M. G., Cicekli, I., & Alpaslan, F. N. (2010). Text summarization of Turkish texts using latent semantic analysis. In *Proceedings of the 23rd International Conference on Computational Linguistics (Coling 2010)*, 869–876.

Pasunuru, R., & Bansal, M. (2018). Multi-reward reinforced summarization with saliency and entailment. *NAACL HLT 2018 - 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies - Proceedings of the Conference*, 2, 646–653.

Pasunuru, R., Guo, H., & Bansal, M. (2017). Towards improving abstractive summarization via entailment generation. *Proceedings of the Workshop on New Frontiers in Summarization*, 27–32.

Paulus, R., Xiong, C., & Socher, R. (2017). A deep reinforced model for abstractive summarization. *ArXiv Preprint ArXiv:1705.04304*.

Rush, A. M., Chopra, S., & Weston, J. (2015). A neural attention model for sentence summarization. In *ACLWeb. Proceedings of the 2015 conference on empirical methods in natural language processing*, 379–389. Association for Computational

Linguistics.

- Sanchez-Gomez, J. M., Vega-Rodríguez, M. A., & Pérez, C. J. (2018). Extractive multi-document text summarization using a multi-objective artificial bee colony optimization approach. *Knowledge-Based Systems*, 159, 1–8.
- Sanchez-Gomez, J. M., Vega-Rodríguez, M. A., & Pérez, C. J. (2020). Experimental analysis of multiple criteria for extractive multi-document text summarization. *Expert Systems with Applications*, 140, 112904.
- See, A., Liu, P. J., & Manning, C. D. (2017). Get to the point: Summarization with pointer-generator networks. *ACL 2017 - 55th Annual Meeting of the Association for Computational Linguistics, Proceedings of the Conference (Long Papers)*, 1, 1073–1083.
- Shi, T., Keneshloo, Y., Ramakrishnan, N., & Reddy, C. K. (2021a). Neural abstractive text summarization with sequence-to-sequence models. *ACM Transactions on Data Science*, 2(1), 1–37.
- Shi, T., Keneshloo, Y., Ramakrishnan, N., & Reddy, C. K. (2021b). Neural Abstractive Text Summarization with Sequence-to-Sequence Models. *ACM/IMS Transactions on Data Science*, 2(1), 1–37. <https://doi.org/10.1145/3419106>
- Sinha, A., Yadav, A., & Gahlot, A. (2018). Extractive text summarization using neural networks. *ArXiv Preprint ArXiv:1802.10137*.
- Speer, R., Chin, J., & Havasi, C. (2017). ConceptNet 5.5: An Open Multilingual Graph of General Knowledge. In *Thirty-first AAAI conference on artificial intelligence* (pp. 4444–4451).
- Subramanian, S., Li, R., Pilault, J., & Pal, C. (2019). On extractive and abstractive neural document summarization with transformer language models. *ArXiv Preprint ArXiv:1909.03186*.
- Sutskever, I., Vinyals, O., & Le, Q. V. (2014). Sequence to sequence learning with neural

- networks. In *Advances in Neural Information Processing Systems* (Vol. 4, pp. 3104–3112). Neural information processing systems foundation.
- Takase, S., Suzuki, J., Okazaki, N., Hirao, T., & Nagata, M. (2016). Neural headline generation on abstract meaning representation. In *Proceedings of the 2016 conference on empirical methods in natural language processing*. (pp. 1054–1059).
- Tan, J., Wan, X., & Xiao, J. (2017). Abstractive document summarization with a graph-based attentional neural model. In *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics*, 1171–1181. <https://doi.org/10.18653/v1/P17-1108>
- Text analysis conference (TAC) dataset*. (n.d.). Retrieved August 9, 2021, from <https://tac.nist.gov/>
- Torun, H., & Inner, A. B. (2018). Detecting similar news by summarizing Turkish news. In *2018 26th IEEE Signal Processing and Communications Applications Conference, SIU, İzmir 2018*, 1–4. Institute of Electrical and Electronics Engineers Inc.
- TS Corpus dataset*. (n.d.). Retrieved August 9, 2021, from <https://ts Corpus.com/>
- TTC-3600: A new benchmark dataset for Turkish text categorization*. (n.d.). Retrieved August 9, 2021, from <https://github.com/denopas/TTC-3600>
- Turkish datasets from Kemik – Yıldız Technical University*. (n.d.). Retrieved August 9, 2021, from <http://www.kemik.yildiz.edu.tr/>
- Turkish paraphrase corpus*. (n.d.). Retrieved August 9, 2021, from <https://osf.io/wp83a/>
- Turkish product reviews*. (n.d.). Retrieved August 9, 2021, from <https://github.com/fthbrmnby/turkish-text-data>
- Turkish sentiment dataset*. (n.d.). Retrieved August 9, 2021, from <http://www.baskent.edu.tr/~msert/research/datasets/SentimentDatasetTR.html>
- Uçkan, T., & Karcı, A. (2020). Extractive multi-document text summarization based on

- graph independent sets. *Egyptian Informatics Journal*, 21(3), 145–157.
- Uzundere, E., Dedja, E., Diri, B., Amasyali, M. F., Fakültesi, E.-E., & Bölümü, B. (2008). Türkçe haber metinleri için otomatik özetleme. In *Akıllı Sistemlerde Yenilikler ve Uygulamaları Sempozyumu, Isparta, Türkiye*, 1–3.
- Van Lierde, H., & Chow, T. W. S. (2019). Learning with fuzzy hypergraphs: A topical approach to query-oriented text summarization. *Information Sciences*, 496, 212–224.
- Wang, L., Yao, J., Tao, Y., Zhong, L., Liu, W., & Du, Q. (2018). A reinforced topic-aware convolutional sequence-to-sequence model for abstractive text summarization. *IJCAI International Joint Conference on Artificial Intelligence, 2018-July*, 4453–4460.
- Webis-Snippet-20 dataset*. (n.d.). Retrieved August 9, 2021, from <https://webis.de/data/webis-snippet-20.html>
- Wu, Y., & Hu, B. (2018). Learning to extract coherent summary via deep reinforcement learning. *Proceedings of the AAAI Conference on Artificial Intelligence*, 32(1).
- Xu, J., & Durrett, G. (2019). Neural extractive text summarization with syntactic compression. *EMNLP-IJCNLP 2019 - 2019 Conference on Empirical Methods in Natural Language Processing and 9th International Joint Conference on Natural Language Processing, Proceedings of the Conference*, 3292–3303.
- Xu, J., Gan, Z., Cheng, Y., & Liu, J. (2019). Discourse-aware neural extractive text summarization. *ArXiv Preprint ArXiv:1910.14142*.
- Yang, M., Wang, X., Lu, Y., Lv, J., Shen, Y., & Li, C. (2020). Plausibility-promoting generative adversarial network for abstractive text summarization with multi-task constraint. *Information Sciences*, 521, 46–61.
- Yousefi-Azar, M., & Hamey, L. (2017). Text summarization using unsupervised deep learning. *Expert Systems with Applications*, 68, 93–105.
- Yu, H., Yue, C., & Wang, C. (2017). News article summarization with attention-based

deep recurrent neural networks. *Technical Report, Stanford University*.

Zeng, W., Luo, W., Fidler, S., & Urtasun, R. (2016). Efficient summarization with read-again and copy mechanism. *ArXiv:1611.03382*.

Zhang, J., Zhao, Y., Saleh, M., & Liu, P. (2020). PEGASUS: Pre-training with extracted gap-sentences for abstractive summarization. In *International Conference on Machine Learning*, 11328–11339. PMLR.

Zhou, Q., Yang, N., Wei, F., & Zhou, M. (2017). Selective encoding for abstractive sentence summarization. *ACL 2017 - 55th Annual Meeting of the Association for Computational Linguistics, Proceedings of the Conference (Long Papers), 1*, 1095–1104.

APPENDICES

Appendix A. Extractive Text Summarization

Author(s)/Year	Summarization Method	Methodology	Used Dataset	Evaluation Metrics
Erkan & Radev (2004)	Extractive	Graph-based	DUC-2003 DUC-2004	ROUGE 1
Barzilay & McKeown (2005)	Extractive	Text-to-text generation technique (MultiGen)	Own Dataset	Precision Recall F Measure
Amancio et al. (2012)	Extractive	Complex networks and syntactic dependency	Brazilian Students Exam Essays	Precision Recall F Measure
Ferreira et al. (2014)	Extractive	Statistics and Linguistic Treatment	DUC 2002	Precision Recall F Measure
Fang et al. (2017)	Extractive	Word Sentence Co-Ranking	DUC 2002	ROUGE 1 ROUGE 2 ROUGE L F Measure
Jain et al. (2017)	Extractive	ANN	DUC 2002	ROUGE 1 ROUGE 2 ROUGE L
Nallapati et al. (2017)	Extractive	Encoder-Decoder RNN	CNN/Daily Mail DUC-2002	ROUGE 1 ROUGE 2 ROUGE L
Yousefi-Azar & Hamey (2017)	Extractive	Encoder-Decoder RNN	BC3 SKE	ROUGE 1 ROUGE 2
Al-Sabahi et al. (2018)	Extractive	RNN with LSTM	CNN/Daily Mail DUC-2002	ROUGE 1 ROUGE 2 ROUGE L
Chen et al. (2018)	Extractive	Encoder-Decoder RNN	CNN/Daily Mail DUC-2002	ROUGE 1 ROUGE 2 ROUGE L
Jadhav & Rajan (2018)	Extractive	Encoder-Decoder RNN	CNN/Daily Mail	ROUGE 1 ROUGE 2 ROUGE L
Khatri et al. (2018)	Abstractive & Extractive	Encoder-Decoder RNN	eBay Human Extracted Snippets	TOKEN SIM ROUGE 1 ROUGE 2 ROUGE L BLEU TOPIC SIM
Sanchez-Gomez et al. (2018)	Extractive	Artificial Bee Colony Optimization	DUC-2002	ROUGE 1 ROUGE 2 ROUGE L

Appendix A. Extractive Text Summarization (Continues)

Author(s)/Year	Summarization Method	Methodology	Used Dataset	Evaluation Metrics
Sinha et al. (2018)	Extractive	2-Layered ANN	DUC 2002	ROUGE 1 ROUGE 2
Wu & Hu (2018)	Extractive	RNN	CNN/Daily Mail	ROUGE 1 ROUGE 2 ROUGE L
Brito et al. (2019)	Extractive	RNN	CNN Articles	ROUGE 1 F1 Score
Keneshloo et al. (2019)	Extractive	Transfer Reinforcement Learning	CNN/Daily Mail	ROUGE 1 ROUGE 2 ROUGE L
Lierde & Chow (2019)	Extractive	Graph Based Approach	DUC05 DUC06 DUC07	ROUGE 2 ROUGE SU4
Liu & Lapata (2019)	Extractive	BERT (Bidirectional Representation from Transformers)	CNN/Daily Mail NYT XSum	ROUGE 1 ROUGE 2 ROUGE L
Madhuri & Kumar (2019)	Extractive	Sentence Ranking	Own dataset	ROUGE 1 ROUGE 2 ROUGE L ROUGE W
Xu & Durrett (2019)	Extractive	Encoder-Decoder RNN	CNN/Daily Mail NYT	ROUGE 1 ROUGE 2 ROUGE L
Xu et al. (2019)	Extractive	Graph based Discourse Aware Neural Network	CNN/Daily Mail NYT	ROUGE 1 ROUGE 2 ROUGE L
Bhargava & Sharma (2020)	Extractive	Deep Neural Networks	Multilingual Single-document Summarization (MSS)	Accuracy Precision F-Measure Recall
Elbarougy et al. (2020)	Extractive	PageRank Algorithm	Essex Arabic Summaries Corpus (EASC)	Precision F-Measure Recall
Ghodratnama et al. (2020)	Extractive	Dynamic Feature Space Mapping kNN	DUC 2002 CNN/DailyMail	ROUGE 1 ROUGE 2 ROUGE L

Appendix A. Extractive Text Summarization (Continues)

Author(s)/Year	Summarization Method	Methodology	Used Dataset	Evaluation Metrics
Hernandez-Castaneda et al. (2020)	Extractive	Genetic Algorithm with Lexical-Semantic Analysis	DUC 2002 TAC11	ROUGE 1 ROUGE 2 ROUGE L Average p-Value
Mohd et al. (2020)	Extractive	Ranking	DUC 2007	ROUGE 1 ROUGE 2 ROUGE L ROUGE SU4
Sanchez-Gomez et al. (2020)	Extractive	Multi-Objective Artificial Bee Colony (MOABC)	DUC 2002	ROUGE 2 ROUGE L
Uçkan & Karcı (2020)	Extractive	Maximum Independent Set	DUC 2002 DUC 2004	ROUGE 1 ROUGE 2 ROUGE L ROUGE W-1 ROUGE W-2 ROUGE SU4

Appendix B. Abstractive Text Summarization

Author(s)/Year	Summarization Method	Methodology	Used Dataset	Evaluation Metric
Lopyrev (2015)	Abstractive	Encoder-Decoder RNN with LSTM	Gigaword	Training and Holdout Loss BLEU
Rush et al. (2015)	Abstractive	Encoder-Decoder RNN	DUC-2004 Gigaword	ROUGE-1 ROUGE-2 ROUGE-L
Chen et al. (2016)	Abstractive	Encoder-Decoder RNN	CNN/Daily Mail LCSTS (Chinese)	ROUGE-1 ROUGE-2 ROUGE-L
Nallapati et al. (2016)	Abstractive	Attentional Encoder-Decoder RNN	Gigaword DUC CNN/Daily Mail	ROUGE-1 ROUGE-2 ROUGE-L
Takase et al. (2016)	Abstractive	AMR Encoder-Decoder RNN	DUC-2004 Gigaword	ROUGE-1 ROUGE-2 ROUGE-L
Zeng et al. (2016)	Abstractive	Read-Again and Copy Mechanism for Encoder-Decoders	Gigaword DUC-2004	ROUGE-1 ROUGE-2 ROUGE-L
Fan et al. (2017)	Abstractive	Encoder-Decoder RNN	CNN/Daily Mail	ROUGE-1 ROUGE-2 ROUGE-L
Kuremoto et al. (2017)	Abstractive	Recurrent Neural Network & Attention-Based Encoder RNN	Two Samples from “The Wall Street Journal”	Training Loss
Li et al. (2017)	Abstractive	Encoder-Decoder RNN	Gigaword DUC-2004 LCSTS (Chinese)	ROUGE-1 ROUGE-2 ROUGE-L ROUGE-SU4
Ling & Rush, (2017)	Abstractive	Coarse-to-Fine Encoder – Decoder RNN	CNN/Daily Mail	ROUGE-1 ROUGE-2 ROUGE-L
Nema et al., (2017)	Abstractive	Encoder-Decoder RNN with Query Based Attention	Debatepedia	ROUGE-1 ROUGE-2 ROUGE-L
Pasunuru et al., (2017)	Abstractive	Encoder-Decoder RNN with LSTM	Gigaword DUC SNLI	ROUGE-1 ROUGE-2 ROUGE-L METEOR BLEU CIDEr-D

Appendix B. Abstractive Text Summarization (Continues).

Author(s)/Year	Summarization Method	Methodology	Used Dataset	Evaluation Metric
Paulus et al., (2017)	Abstractive	Encoder-Decoder RNN with Reinforcement Learning	CNN/Daily Mail New York Times	ROUGE-1 ROUGE-2 ROUGE-L
See et al. (2017)	Abstractive	Encoder-Decoder RNN	CNN/Daily Mail	ROUGE-1 ROUGE-2 ROUGE-L METEOR
Tan et al. (2017)	Abstractive	Encoder-Decoder RNN	CNN/Daily Mail	ROUGE-1 ROUGE-2 ROUGE-L
Zhou et al. (2017)	Abstractive	Encoder-Decoder RNN	Gigaword DUC 2004 MSR	ROUGE-1 ROUGE-2 ROUGE-L
André et al. (2018)	Abstractive	A Diverse Beam Search Approach	CNN/Daily Mail	ROUGE-1 ROUGE-2 ROUGE-L
Celikyilmaz et al. (2018)	Abstractive	Encoder-Decoder RNN	CNN/Daily Mail NYT	ROUGE-1 ROUGE-2 ROUGE-L
Chen & Bansal (2018)	Abstractive	Encoder-Decoder RNN with Reinforced Learning	CNN/Daily Mail DUC-2002	ROUGE-1 ROUGE-2 ROUGE-L METEOR
Cohan et al. (2018)	Abstractive	Encoder-Decoder RNN with Attention Model	arXiv PubMed	ROUGE-1 ROUGE-2 ROUGE-3 ROUGE-L
Gehrmann et al. (2018)	Abstractive	Encoder-Decoder RNN	CNN/Daily Mail NYT	ROUGE-1 ROUGE-2 ROUGE-L
Hsu et al. (2018)	Abstractive & Extractive	Encoder-Decoder RNN	CNN/Daily Mail	ROUGE-1 ROUGE-2 ROUGE-L
Jiang & Bansal (2018)	Abstractive	Encoder-Decoder RNN	CNN / Daily Mail DUC-2002	ROUGE-1 ROUGE-2 ROUGE-L
Kryściński et al. (2018)	Abstractive	Encoder-Decoder RNN	CNN/Daily Mail	ROUGE-1 ROUGE-2 ROUGE-L
Li et al. (2018)	Abstractive	Encoder-Decoder RNN	DUC-2004 LCSTS (Chinese) Gigaword	ROUGE-1 ROUGE-2 ROUGE-L

Appendix B. Abstractive Text Summarization (Continues).

Author(s)/Year	Summarization Method	Methodology	Used Dataset	Evaluation Metric
Li et al. (2018)	Abstractive	Encoder-Decoder RNN	CNN/Daily Mail	ROUGE-1 ROUGE-2 ROUGE-L
Lin et al. (2018)	Abstractive	Encoder-Decoder RNN	Gigaword LCSTS (Chinese)	ROUGE-1 ROUGE-2 ROUGE-L
Liu et al.(2018)	Abstractive	Reinforcement Learning	CNN/Daily Mail	ROUGE-1 ROUGE-2 ROUGE-L
Pasunuru & Bansal (2018)	Abstractive	Encoder-Decoder RNN with Reinforcement Learning	CNN/Daily Mail DUC-2002 SNLI Multi-NLI	ROUGE-1 ROUGE-2 ROUGE-L METEOR
Wang et al. (2018)	Abstractive	Convolutional Sequence-to-Sequence NN	DUC-2004 dataset & Gigaword dataset	ROUGE-1 ROUGE-2 ROUGE-L
Bae et al. (2019)	Abstractive	Reinforcement Learning	CNN/Daily Mail New York Times	ROUGE-1 ROUGE-2 ROUGE-L
Liu et al. (2019)	Abstractive	Encoder-Decoder RNN with Auto Encoding	Internal dataset collected through Amazon Mechanical Turk	ROUGE-1 ROUGE-L
Subramanian et al. (2019)	Abstractive	Hierarchical Encoder-Decoder Sentence Pointer	ArXiv PubMed Newsroom BigPatent	ROUGE-1 ROUGE-2 ROUGE-3 ROUGE-L
Gunel et al., (2020)	Abstractive	Encoder-Decoder RNN	CNN/Daily Mail	ROUGE-1 ROUGE-2 ROUGE-L
Liang et al. (2020)	Abstractive	Encoder-Decoder RNN	LCSTS	ROUGE-1 ROUGE-2 ROUGE-L
Liao et al. (2020)	Abstractive	Encoder-Decoder RNN	CNN/Daily Mail	ROUGE-1 ROUGE-2 ROUGE-L
Moirangthem & Lee (2020)	Abstractive	Encoder-Decoder RNN with Temporal Hierachy	CNN/Daily Mail	ROUGE-1 ROUGE-2 ROUGE-L
Yang et al., (2020)	Abstractive	Plausibility-promoting Generative Adversarial Network	CNN/Daily Mail Gigaword	ROUGE-1 ROUGE-2 ROUGE-L
Zhang et al. (2020)	Abstractive	Transformer-based encoder-decoder model with self supervised objective	CNN/Daily Mail Xsum Reddit TIFU	ROUGE-1 ROUGE-2 ROUGE-L

Appendix C. Text Summarization in Turkish

Author(s)/Year	Summarization Method	Methodology Applied	Used Dataset	Evaluation Metrics
Altan (2004)	Extractive	Statistical Analysis	50 different articles about economics	X
Karakaya & Güvenir (2004)	Extractive	Statistical Analysis	11 different articles related economics crisis in Turkey	X
Uzundere et al. (2008)	Extractive	Statistical Analysis	10 different articles	Standard Deviation
Kutlu et al. (2010)	Extractive	Generic Text Summarization	120 newspaper articles and 100 Turkish journal articles	ROUGE-1 ROUGE-2 ROUGE-L ROUGE-W
Ozsoy et al. (2010)	Extractive	Latent Semantic Analysis	Two different data sets of scientific articles in Turkish (including articles from various areas)	ROUGE-L F-measure score
Güran et al. (2013)	Extractive	A hybrid summarization method	110 news documents and 3 analysts	ROUGE-1
Hatipoğlu & Omurca (2015)	Extractive	A hybrid summarization	10 different documents from Wikipedia	Selection Ratio
Hatipoglu & Omurca (2016)	Extractive	A hybrid summarization	10 different documents from Wikipedia	Average precision and breakeven point
Demirci et al. (2017)	Extractive	Latent Semantic Analysis	106 selected news	ROUGE-1
Torun & Inner (2018)	Extractive	TF-IDF	12000 News	Average Calculation
Doğan & Kaya (2019)	Extractive	Latent Semantic Analysis	30000 Tweets	F-measure score
Hark et al. (2019)	Extractive	Graph Entropy	DUC-2002	ROUGE 1 ROUGE-2 ROUGE-3 ROUGE-4
Karakoc & Yılmaz (2019)	Abstractive	RNN	“Kemik haber” dataset	ROUGE-1

Appendix D. Five Examples of the Collected Data from the TR-News-Sum Dataset

1st Example.
Main Text
Twitter, çalışanlarına korona virüsün yayılmasını durdurmak için evden çalışmalarını söyledi. Bir blog yazısında, sosyal medya'nın devi Hong Kong, Japonya ve Güney Kore'deki personelin uzaktan çalışmasının zorunlu olduğunu belirtti. Şirket ayrıca tüm dünyadaki 5 bin çalışanının Twitter ofislerine gitmemesini "şiddetle teşvik ettiğini" söyledi. Ayrıca sonraki günlerde Twitter yönetimi çalışanlarının gerekli olmayan iş seyahati ve etkinliklerini askıya aldığı şirket içinde duyurdu. Şirket, Austin ve, Teksas'taki bu ay gerçekleşecekın South by Southwest medya konferansından çekildiğini zaten duyurmuştu. Twitter'ın insan kaynakları başkanı Jennifer Christie: "Amacımız, Covid-19 korona virüsünün bizim için ve etrafımızdaki dünyanın yayılma olasılığını azaltmak." dedi. Christie mesajda ayrıca Twitter'ın bir süredir evden çalışmanın yollarını geliştirdiğini de vurguladı. "Bu bizim için büyük bir değişiklik olsa da, giderek daha uzak olan daha dağıtılmış bir iş gücüne doğru ilerliyoruz. Küresel bir hizmetiz ve biz herhangi bir yerde, Twitter'da çalışabilmeyi sağlama konusunda kararlıyız." dedi. Twitter'ın kurucusu Jack Dorsey uzun zamandır uzaktan çalışma iş modelini destekledi ve Kasım ayında bu yılın altı aya kadar olan kısmını Afrika'da geçirmeyi planladığını açıkladı. Twitter'ın bu hareketi, virüs bölgeyi etkisi altına aldığı için için Asya'daki birçok şirket tarafından uygulanan önlemlere benziyor, ancak Twitter'ın salgına yanıtı çoğu büyük Amerikan işletmesinden daha da ileri gitmiş durumda. Facebook ve Google dahil olmak üzere diğer önde gelen teknoloji şirketleri ABD'deki konferansları erteledi veya iptal etti. Facebook da Twitter'a katılarak, Southwest medya konferansından çekildi. Google'ın Dublin'de bulunan Avrupa genel merkezindeki personel, İrlanda'daki potansiyel bir salgın için hazırlığını test ederken Salı günü evden çalışacak, ancak 8 bin çalışanın çoğunun çarşamba günü masalarına dönmesi bekleniyor. Aynı zamanda, telekom operatörü A&T ve bankacılık devi Citigroup da dahil olmak üzere çoğu şirket, özellikle Asya'ya uluslararası seyahatlerini kısıtladı.
Summaried Text
Twitter, çalışanlarına koronavirüsün yayılmasını durdurmak için evden çalışmalarını söyledi. Twitter'ın kurucusu Jack Dorsey uzun zamandır uzaktan çalışma iş modelini desteklediğini ve Kasım ayında bu yılın altı aya kadar olan kısmını Afrika'da geçirmeyi planladığını açıkladı.
Category
Genel
Headline
Twitter Yönetiminden Korona Virüs Kararı

Appendix D. Five Examples of the Collected Data from the TR-News-Sum Dataset (Continues)

2nd Example.
Main Text
Milliyet Ege Spor Müdürü Mehmet Demirtaş, Dünyevi Haber muhabiri Samet Adilak'a gazetecilik mesleğinin sorunlarını, geleceğini ve gündemini anlattı. Yeni medya ve sosyal medyanın gösterdiği gelişimle haberciliğin bu alanlara eksen kaydırıldığını belirten Demirtaş, "Güncelliği yakalayan haberciler, günümüzde daha geniş kitlelere ulaşmak için yenilikleri yakından takip etmeliler" dedi. Gazetecilik kariyerine Gözlem Gazetesi'nde başlayan daha sonra Gazete Ege'de devam eden ve son 7 yıldır da Milliyet Ege'de çalışan Mehmet Demirtaş, Dünyevi Muhabiri Samet Adilak'ın sorularını samimiyetle yanıtladı. Gazetecilik mesleğinin sorunlarına, gündemine, geleceğine dair açıklamalarda bulunan Demirtaş, "Yetişen yeni nesil teknoloji çağının gereklilikleri ile büyüyor. Bu çağda habere ulaşmak da bilgiye ulaşmak da her zaman kolay olsun isteniyor. Yeni nesil okumak alışkanlığından çok görmek ve izlemek üzerine yoğunlaşıyor. Habercilerin de bu çağın gerekliliklerini yakından takip ederek, doğru konum alması gerekiyor" diye konuştu. * Neden gazetecilik? Ben anne ve babası gazeteci olan, haberciliğin içine doğmuş bir çocuktum. Hep gazete, ajans ve haber kavramlarını özümseyerek büyüdüm. Gençliğimde harçlığımı kazanmak için başka işlerde çalışsam da hayat beni yine bir şekilde gazetecilikle buluşturdu. Uzak olmadığım, tanıdığım kavramların içinde buldum kendimi. Mesleğin başında sayfa sekreterliği için başladım. Usta ellerden eğitim gördüm. Büyük isimlerle çalıştım. Kendimi sayfa sekreterliğinde geliştirmek için çok çalıştım. Bu çabamın meyvesini de hep topladım. Gazetecilik her ne kadar zor, yorucu, stresli ve yarınının garanti olmadığı bir meslek olsa da benim için bir tutku oldu. Sayfa sekreterliğinden, spor müdürlüğüne adım attığımda çabalarımın ve emeğimin yanıtsız kalmadığını gördüm. Şimdi ise tutkunu olduğum bu meslekte olgunluk çağıma geldi. * Çevrenizin, mesleğiniz adına algıları ve tutumları neler oluyor? İnsanlar saygın ve prestijli bir işin parçası olduğumuzu düşünüyorlar. Bence de bu durum böyle. Gazetecilik bana kalırsa da çok saygın ve önemli bir meslek. Her insanın yapamayacağı, her haber okuyanın haber yazamayacağı, her program bilenin sayfa sekreterliği yapamayacağı bir gerçek. İnsanlar 'Herkesten gazeteci olur' algısını yıkmalı. Bu durumun böyle olmadığını bilen insanlar bizim mesleğimize en azından biraz daha olsun bizim gözümüzden bakabiliyor. İşte o zaman bu işin ne kadar keyifli bir iş olduğunu hissediyor insan. Başka erkek kardeşim olmak üzere yakın çevrem çöğü gazeteci. Bbu sebepten genellikle onlarla aynı konular üzerinde durabiliyoruz. * Gazetecilik mesleğinin dünü, bugünü ve yarını hakkında neler düşünüyorsunuz? Gazetecilik mesleğinin dününe baktığımızda yoğunluğu daha çok, ilgisi çok, merak duyulan ve çok kazandıran bir meslek olarak görülmüştür. Geçmişte bu meslekte iyi paralar kazanılıyormuş. Gazeteciler haklarını bilen, sorumlu haberciliğe önem gösteren ve sadece kendisini değil aynı zamanda çevresini de gözetken bireylemiş. Günümüzde bu tarz kavramlarının çoğunun yozlaştığını görüyoruz. Öte yandan da günümüzde gazeteciliğin herkesçe yapılabilecek olan bir meslek olduğunu, elinde akıllı telefonu olan herkesin potansiyel bir gazeteci olduğunu düşünmek gibi büyük ve yanlış bir algı var. Gazetecilik mesleğinde günümüzde büyük paralar kazanma hayali kuran kim varsa hayal kırıklığına uğrar. Bunu açıkça söyleyebilirim. Bu meslek biraz gönül işidir. Beklediği paraları kazanamayan çoğu yeni gazeteci adayı mesleği için koşturmayı bırakıp başka alanlara yönelebiliyor. Onları da hoş görüyorum. Yaşam zor ve pahalı. Bu mesleğin yarısını ele almak gerçekten güç. Bildiğim tek bir şey var. O da gazetecilikte çağın gerekliliğini yakalayamayan kim varsa tarih olup unutulacaktır. * Bu mesleği yapma hayalleri kuran ve Gazetecilik bölümünden mezun olacak olan yeni gazeteci adaylarına tavsiyeleriniz neler? İnsan hangi işi yaparsa yapsın sevmesi gerekiyor. Sevmeden yapılan her iş, çalıştığın her gün için o kişiye zulüm olur. Öncelikle gazetecilik mesleğine saygı duymalı ve bu mesleği sevmeliler. Yabancı dil konuşundan muhakkak yatırım yapmalılar. Bilgisayar dilini iyi öğrenip program kullanmayı bilmeliler. Halkın içinde olmalı, insanların sorunlarından dertlerinden kopmamalılar. Makam ve koltuk sevdalarını, kariyerlerinin önlerinde tutmamalılar. Doğru ve gerçek haberciliğe, gazetecilik etiğine önem göstermeliler. Yılmamalılar. Bu mesleği gerçekten yapmak isteyen kim varsa teknolojiyi, çağı, yeni medyayı doğru okumalı. Doğru anlamalı ve uygulamalılar.
Summaried Text

Röportaj, meslek hayatında 10 yılı geride bırakan Milliyet Ege Spor Müdürü Mehmet Demirtaş'ın meslek hayatı ve 'Gazetecilik' konu başlıkları altında gelişmiştir.
Category
Genel
Headline
"Gazetecilik gönül işidir"



Appendix D. Five Examples of the Collected Data from the TR-News-Sum Dataset (Continues)

3rd Example.
Main Text
Şiddetin en çok görüldüğü sektör, sağlık sektörü. Tüm sağlık çalışanları kadar, hasta ve hasta yakınları da tehdit altında. Kendini güvende hissedemeyen çalışanlar işlerine yeterli özverişi sağlamakta güçlük çekiyor. Sağlık çalışanlarının %81,9'u, görevlerini yerine getirirken şiddetle karşı karşıya gelme konusunda endişe yaşadıklarını belirtiyor. Dolayısıyla hastalar da bu durumdan olumsuz etkileniyor. 2019 verilerine göre Türkiye'de kayıt altına alınmış sağlıkçıya şiddet (beyaz kod) sayısı 89 bin 172 (!) Veriler bu şiddetin kaynağı olarak büyük oranda hasta yakınlarını işaret ediyor. Mevcut güvenlik önlemlerini; güvenlik görevlileri, beyaz kod, çalışan hakları birimi, kamera sistemleri olarak tanımlayan çalışanlar, bu önlemlerin yeterli olmadığını ya da işlevselliklerinin tam anlamıyla bulunmadığını belirtiyor. Elde ettiğimiz verilere göre hastanelerin sahiplik yapısına göre kayıt altına alınan şiddet oranları değişiyor. Şiddet mağduru sağlık çalışanının, saldırgandan şikayetçi olup hukuki süreci başlatmasına beyaz kod adı veriliyor. Sadece beyaz kod veren çalışanların şikayetleri 'sağlıkta şiddet verisi' olarak sayılıyor. Beyaz kod sayısı; devlet hastaneleri ve üniversite hastanelerinde oldukça yüksek gibi görünse de esasen özel hastanelerde kayıt altına alınmadığı için az ya da hiç yokmuş gibi görünmekte. Can güvenliğinden endişe duyan sağlık çalışanlarına onların neler düşündüğünü ve neler hissettiklerini sorduk. Devlet hastanesi, özel hastane ve üniversite hastanesi çalışanlarıyla konuştuk: Merve Çalışır "Ben bir devlet hastanesinin acil servisinde hemşire olarak görev yapmaktayım. Elbette kendimi güvende hissetmiyorum. Gerek hasta gerekse hasta yakınlarının her an bana zarar verebileceğini düşünüyorum. Daha önce beyaz kod da verdim ama pişman oldum çünkü dava süreci o kadar uzun ve yorucu ki beyaz kod vermeye bile değmiyor diyorsun. Zaten psikolojik ve sözel rahatsız edici şeylere alıştık ama fiziksel şiddette bile kendimi koruyayım ve mevzu daha fazla uzamasın diye olaya bakıyorsun. Şu an devam eden bir davam zaten var ve iki yıldır hiçbir sonuç alamadım." Sağlık kurumlarında şiddetin en fazla görüldüğü yer acil servisler. Poliklinik sırası beklemek istemeyen hastalarla amacı dışında hizmet veren aciller oldukça kalabalık. Yoğunluk hem sağlık çalışanları hem de hastalar üzerinde stres oluşturuyor. Buna ek olarak doktorlara uygulanan performans sistemi yüzünden gerekli gereksiz tüm hastalardan istenen MR, BT, kan tahlilleri vb. durumlarla hastaların, hastanede kalış süreleri uzuyor. Sağlıkta uygulanan 'performans uygulamaları', hekimlerin insan yaşamına, sağlığına özen göstermelerini engellemekte. Bu uygulamadan dolayı hekimler hastaların sağlığı konusunda düşünmekten çok ne kadar çok hasta bakarım derdine giriyor. Dolayısıyla hastalara verilen sağlık hizmetinin niteliğini düşürürken hekimler arasında da niceliksel rekabete sebep oluyor. Performans sistemi hekimler arasında etik dışı yarışmayı tetikleyip sağlık harcamalarını ölçsüz ve gereksiz oranda artırmakta. Aşırı finansal teşvikler sadece etik değerleri bozmayıp ayrıca hukuki sorunlar da oluşturuyor. Uygulanan performans dayalı ücret rejimi ekip dayanışmasını parçalayan, hekimi hekime rakip kılan, hekimi ekibin üyesi olan diğer sağlık emekçilerine yabancılaştıran, insanı ilişkileri zedeleyen bir işlev görmektedir. Performansa dayalı çalışma sisteminde hekimlerde izin kullanmamaya bağlı aşırı stres ortaya çıkarıyor. Hekimler; değişen oranlarda endikasyonsuz müdahalelerin, etik olmayan uygulamaların malpraktis ve komplikasyon sayısının arttığı, verilen sağlık hizmetinin niteliğinin ve hasta başına düşen muayene süresinin düştüğünü ifade ediyor. Hasta Değil Müşteri Burcu Bahçekapılı "Ben özel hastanede hemşire olarak çalışıyorum. Daha önce hiç şiddete uğramadım ama zaten öyle bir durum söz konusu bile değil. Beyaz kod hakkında eğitim aldık evet ama bunu kullanmak aklımızın ucundan bile geçmez çünkü ucu bize dokunur. Hemşire ile hasta bakıcı arasında bir pozisyonda gibiyiz. Hastalara sağlık hizmeti vermenin ötesinde sanki bir otel hizmeti de veriyoruz. Onları hasta olarak değil müşteri olarak görüyoruz. Hastanenin genel politikası bu yönde. Arabalarına kadar uğurluyoruz. Ne bir hakkımız ne de saygınlığımız mevcut..." Hastanelerin sahiplik yapısı da çalışan haklarını doğrudan etkilemekte. Devlet hastanelerinde ya da üniversite hastanelerinde haklarını biraz olsun arayabilen, koruyabilen çalışanlar, özel hastanelerde herhangi bir hak arayışında bulunamıyor. Ekonomi-politik meseleler yüzünden hastanın müşteriye dönüşme süreci insan hakları açısından da oldukça düşündürücü. Ufuk Taşçı " Ben devlet hastanesinde uzman doktor olarak görev yapıyorum. Kendimi 'kısmi' güvende hissediyorum.

Sözel şiddet çok fazla yaşanıyor. Fiziksel şiddete ise 1 kez maruz kaldım. Dava 2 yıldır yargı sürecinde.” Dava süreçlerinin uzun sürmesi mağdurların şikayetçi olma oranını da olumsuz etkiliyor. İnsanlar yıllar boyunca mahkeme koridorlarında olmak istemiyor. Her ne kadar bakanlık avukat desteği sağlasa da görünen o ki bu da yeterli gözükmüyor. Feyyaz Döner “Üniversite hastanesinde çalışıyorum. Çalışma şartları çok zorlu. Hem çalışma saatleri hem iş yoğunluğu üstüne bir de kalabalık hasta yığınları derken durum içinden çıkılmaz bir hal alıyor. Psikolojik şiddete de fiziksel şiddete de maruz kaldım fakat beyaz kod vermedim. Çünkü başıma gelecek süreci biliyorum ve uğraşmak istemedim. Kendi yöntemlerimle (gülüyor) hallettim.” Dava sürecinden kaçınan kimi çalışanlar ‘kendi yöntemleri’ ile hallettiklerini belirtse de kadın çalışanlar için bu durum pek de mümkün değil. Sağlık sektöründe çalışan kadın personel kariyeri boyunca yüzde 90,50 ihtimalle en az bir kere şiddete maruz kalıyor. Toplumsal cinsiyet rollerinden ötürü kadınların şiddete maruz kalma oranı daha yüksek. Burak Fidandan “Yönetim ve çalışma psikolojisi alanında yüksek lisans yapıyorum. Sağlıkta uygulanan şiddetin aslında pek çok nedeni var. Şiddet kültürü, hasta ve yakınlarına ait faktörler ve sistemle ilgili sorunlar bunun temel nedenleri. Paternalistik tutumun da etkisi var tabii. Şiddete prim vermeyen ahlaki ve hukuki kültürümüz maalesef yok. Televizyonlarda, medyada şiddeti meşrulaştıracak her şey gösteriliyor. Bunların da bu konu hakkında yardımcı olmadığını belirtmek lazım. Şiddet kültürünün yaygınlaşmasında aslında medya başı çekiyor. Şiddet haberlerinde kullanılan kutuplaştırıcı dil hasta ve sağlık çalışanı arasında ön yargı oluşturmakta. “Dayak yedi” , “Reçete yazmadığı için...”, “Sıra beklemek istemedi...” vb. kalıplarla yapılan haberlerin de şiddete sevk etme payı hiç de az değil. Medya halkı mı yansıtıyor yoksa medya mı halkı bilinmez ancak bu dilemmanın bir şekilde çözülmesi gerekiyor. Yaşam boyu eğitim modeliyle sağlık okuryazarlığı artırılmalı, devlet ve özel sektör işbirliğiyle halkın sağlık hususunda daha bilinçli hareket etmesi sağlanmalıdır. Sonuç olarak sağlıkta şiddet problemi aslında kişiye özel gibi gözüke de tüm toplumun problemi. Bunu önleyecek çalışmalar için devlet daha fazla politika oluşturmalı, bu konuda uzman görüşleri değerlendirilmeli ve uygulanmalı.

Summaried Text

Sağlıkta şiddetin kodu olan “beyaz kod” uygulaması üniversite hastanelerinde, devlet hastanelerinde ve özel hastanelerde değişiklik gösteriyor. Çalışanlar, çalıştıkları hastanenin sahiplik yapısı ya da hukuki sürecin oldukça uzun olması nedeniyle beyaz koddan kaçındıklarını belirtti. Eldeki sayısal veriler net bir sonuç veremese de rakamlar bir hayli yüksek. Şiddet mağdurlarının ezici çoğunluğunu kadınlar oluşturuyor. Sağlık çalışanlarının büyük çoğunluğu canlarının güvenliğinden endişe ediyor

Category

Genel

Headline

Sağlıkta Şiddet Değil Sağlığa Şiddet

Appendix D. Five Examples of the Collected Data from the TR-News-Sum Dataset (Continues)

4th Example.
Main Text
Türkiye de ötelenen bir müzik dalını trendlere sokan, yaptığı müzik için tutuklanan, sokaklarda büyüyen ve bugün Türkiye'nin yükselen yıldızı haline gelen Serca İpekçioğlu namı diğer "Ezhel" ile sohbet ettik. -Merhaba Ezhel, İzmir'e hoşgeldin. Sahnede hakikaten seyirciyle çok samimi bir etkileşimin var. Şu anda çok profesyonel bir şey yapıyorsun ama içindeki amatör ruh ölmemiş. Bbu böyle devam edecek mi yoksa korkuyor musun bir gün içindeki o ruh bozulabilir diye korkuyor musun? Çok büyük bir kitleye sahipsın ve gittikçe artıyor. Hiç tahmin etmiş miydin böyle bir şey olacağını? Selam, hoş buldum. Aslında hiç tahmin etmemiştim. Yıllardır müzikle iç içe olduğum için içte olan bir şey gibi geliyordu bana. Sonradan kaybedebileceğin ya da kazanabileceğin bir şey değil. Hep olan. Tabii yorucu bir şey. 287 yaşındayım şu an için bana da rahat geliyor. Bu tempoyu kaldırabiliyorum en azından ama ileride kendime de vakit ayırmak istiyorum bir noktadan sonra. Arkama çekilip şöyle sadece kafamı dinlemek için müzik yapmak, sadece kendim için belli bir kaygı gütmekten, endüstriyel bir kaygı gütmekten, maddi bir kaygı gütmekten bunu yapmak isterim. Ama müzik insanın içinde olduktan sonra istediği yaşa gelsin yine kalıyor, o doğal bir şey. Çıkmak istiyor zaten sürekli. -Özgür olduğunu düşünüyor musun? Tabii kesinlikle. Çünkü bir plak şirketiyle çalışmıyorum. Zaten en başta aldığımız bir karardı bu. Çünkü şu an biraz daha oyun değişiyor ama ben albümü yaparken Mayıs 2017'ye2017'de mayısa baktığın zaman o zaman hiçbir şey yok. Ne "Ezhel" var ne "Müptezhel" var. O zaman plak şirketlerinin sana bakış açısı tabii ki çok düz oluyor. Umursamıyorlar bile diyebilirim. Mesela o noktada çok komik fiyatlar teklif ediyorlar. Albümü haklarını almak için çok komik rakamlar vardı teklif edilen. 15 bin lira gibi. Albüm için sadece. Bana sadece 15 bin lira verecekler albüm tamamen onların. Geri kalan her şey de onlara ait. Komikti, zamanla albüm çıktıktan sonra, biraz ses yaptıktan sonra ben yine şans verdim. Bakalım ne diyor müzik şirketleri diye. Bir iki görüşmeye de gittim ama vahim geldi yine durum bana. Çünkü adam vurduğunu indiriyor. Seni sindiriyor, istediğini alıyor bir şekilde. Ama bu şirketin mantalitesiyle alakalı. Her şirket kısıtlamıyor. Bazısı gerçekten sanatçısı için uygun çözümler arıyor. Belli yollar arıyor üstüne gidiyor. Sırf maddiyat değil. Ben demiyorum ki şirket kazanmasın. Her şey benim olsun demiyorum. Herkesin belli bir işi var. Herkesin belli bir dağılıma ihtiyacı var. Ama ortada belli bir sömürü durumu var müzik piyasasında. Şirket tabanlıbazlı bir sömürü hali var. Zaten biz hep internet sayesinde yaptık bugüne kadar yaptıklarımızı. Youtube bir ara yükselirken biz milyonluk klibi gördüğümüzde şaşırıyorduk o zamanlar. Şimdi tabii bambaşka tabii ama internet zaten hep hayatımızda vardı. Kendimiz de kazanabilecekken, böyle bir şey mümkünken kendimiz yaptık. Rapçilerin çoğu da beat (altyapı) yapmayı bilir. Çünkü yapan kimseyi bulamaz. Çoğu video düzenlemeyi de öğrenir. Çünkü kimsen yok. Biraz seni kendin yapmaya iten bir şey bizim bulunduğumuz müzik. Biz tabii şirketlerle çalışıyoruz ama özgürüz. Şarkılarımızı koymak için belli araçlar lazım ama dediğim gibi tamamen bizim avantajımız da olan daha sanatçı bazlı bir haldeyiz. -Ezhel 1995'te çıksaydı nasıl olurdu? Çok güzel bir soru. Öncelikle başka bir tarz olurdu muhtemelen. O gün duyduğumuz şekilde olmazdı muhtemelen. Daha yine çağın müziği olurdu. 95'in müziği olurdu ama yine hemen hemen politikayla ilgili yapardım. 90'lar çok sert dönemlermiş. Bu kadar olur muydu bilmiyorum. 2018'de bile insanın sözlerden başı belaya girebiliyor o zamanlar. Tabii ince bir çizgi. Siyasetin ve magazin çok ince dengeler içinde olduğu zamanlar. 95'te 4 yaşındaydım çok hatırlamıyorum ama kalan zamanlardan hatırladığım şeyler oluyor. Bu kadar açık olur muydum bilmiyorum. Belki yine olurdu ama politik ve siyasi söylemlerim daha mı hafif olurdu daha mı ağır olurdu? Hayatımda duyduğum en güzel sorulardan biriydi. -Ezhel'in ben de bir mesaj vermeliyim, toplumun bir şekilde derdini sıkıntısını underground olarak bir şekilde yansıtmalıyım diyor musun? Diyorum. Mesela albüm ilk çıktığında beni üzen bir şeydi. Şarkım var bir tane. Dediğin gibi "Benim derdim" hatta şarkının adı ve saçma dertlerden tekrar dünyanın ve toplumun dertlerini anlatmaya çalıştığım bir şarkı ama insanların aklında kalan tek şey "sosis ve salam" oluyor. Sanki başka hiçbir şey yokmuş gibi. O yüzden biraz insanlar için yapmayacaksın. Onu anladım. Kendin için kendini tatmin etmek için yapacaksın. Toplum da aslında sadece almak istediğini alıyor senden ama sen onlara vereceğin şeyi geniş tutarsan onu da alır. Kaygı güdüyor muyum, içimde böyle bir şey yapmalı mıyım diye, hayır,

o an için özgürce içindenm geliyorsa geliyor. Sürekli toplumsal bir meseleye gitmek de bana sahte geliyor. Sürekli toplumsal bir meseleye gitmek de bana sahte geliyor. Sanki bu bir kulvarmış ve ben tamamen buradan yürüyormuşum gibi. -Ezhez daha underground bir kitleye mi hitap ediyor? Rap müziğin eskiye nazaran sınırları genişledi mi? Ben genişlediğini görebiliyorum. Bundan da aşırı mutluyum. Çünkü çok uğraştık bunun için. Sadece kendimiz için de değil. Ben her zaman Türkçe Rap'in bir endüstriye dönüşmesini de düşündüm. Bireysel bir başarı elde edersen yalnızca bireysel başarı elde etmiş olursun. Bundan ben de kazanamayacaktım. Gelebileceğim noktanın çok altında kalacaktım. Benim kazanmam için Türkçe Rap'in de kazanması lazım. Ben hep böyle gördüm ve düşündüm. Şu anda geldiğimiz konumdan çok mutluyum. Rap şarkılarımız dizilerde, magazinlerde çıkmaya başladı. Sürekli bizden bahsediyorlar, var tabii artıları eksileri ama hala underground insanlar da var. Ben zaten bu başarıyı da genel sağlam rap dinleyicisine borçluyum. Kabuk rap dinleyicisine borçluyum. Çünkü ben albümü yaptığımda o insanların aşırı dinlemesi ve paylaşımı sayesinde belli bir şekilde Spotify'da viral listesine girdim. bütün Türkiye'ye alakasız insanlara bile yayılsa ben şu an geldiğim noktayı underground rap'i takip eden kemik kitleye borçluyum ama yayılsın. Bu da çok güzel bir şey. Bazı insanlar çok konservatif bakıyor bu duruma. Rap şöyle olmasın, aman Ezhez de rap mi popçu diye. Bunlar da hiç umurunda değil, hiç gıcınmıyorum. İstediklerini desinler. Ben cayır cayır rap yaptığımı bildiğim için bana koymuyor. Haberlerde 'popçu' yazmışlar, yüzeysellik bizim ülkemizde kültür. Bir şeylere sadece kabuğundan bakmak sadece görüntüsünü almak, bununla yargılamak bir kültür. Ben hiç gıcınmıyorum, aşırı mutluyum. Çünkü başardık. Ben asla popçu olmadım. Ben Türkiye'de asla pop dediğimiz müziği yapmadım. Yaptığımız rap, pop olduysa yapacak bir şeyimiz yok. -Bir rap sanatçısı ilk defa bu kadar büyük konser salonlarını dolduruyor. Bbu kadar geniş bir kitle tarafından takip edilmek seni korkutuyor mu? İlk başlarda oldu. Zor tabii ne yaparsan yap, sosyal hayatın ister istemez azalıyor. Ben sokağı çok seviyorum. Sokakta gezmek bir yerden bir yere gitmek yürümek... Bunu kötülemiyorum. İnsanların sevgisi müthiş bir şey. İnsanı dinç tutuyor, mutlu ediyor bazen ama bir yandan da yorucu olabiliyor. Herkesin talebine cevap vermek zorundasın. Kimseyi kırmaman lazım. bu açılarından zor. Ben kendi somut hayatımdan ürktüyorum. Yoksa genel olarak keşke bütün dünya bilse de müziğimizi yapsak, paylaşsak, dinlenebilsek. Sadece böyle küçük sosyal problemler var. Bu da herkesin başına gelebilir. -90'ların sonunda 2000'lerin başına doğru rock ve metal müzik daha çok ön plandaydı. devir değişti mi? yükselen yeni yıldız rap mi? Aynen,. hatta Türkiye'de rock, bizim rapin yaşadığı çıkışı yaşamıştı Türkiye'de rock. İnanılmaz bir dönemdi. Ben bayılırdım o zaman, üzgünüm baya. Ben de sağlam rock müzik dinleyicisiydim. Türkçe rock'ın Türkçe rap gibi çıktığı bir dönem vardı. İnanılmaz. Gruplar sürekli festivallerde; Çilekeş Direc-t, Kurban, Duman... Sürekli yeni bir şeyler çıkarıyorlar, deneyimler Türkiye için. Açıkçası Türkiye'de yapılan herhangi bir şeyle gurur duyuyorum zaten. Ama üzüliyorum rock müziğin geldiği noktaya. Ben şöyle düşünüyorum, bu bir devir daim. Rapin aslında şöyle bir gücü var: rap uyum sağlayabilen bir müzik. Biraz su gibi. bulunduğu kabın şeklini rock'a göre daha rahat alabilen bir tür. Çünkü rap'i bir enstrüman gibi düşün altyapısıyla. Bunu ritim olan herhangi bir şeyin üstüne rap yapabilirsin,. bariyerleri biraz daha sınırsız. Biz bugün Ttrap kullanıyoruz, auto-tone'la müzik yapıyoruz. Bbu demek değil hep böyle gidecek. 2010'ların başında rap, house müzikle bir ara çok iç içe girmişti. Rr&Bb, soul müzik, her zaman kendine bir yer ve kalıp bulabilen bir tarz. Şu an için yükselen yeni yıldız rap, Türkçe rock zamanı için de şöyle bir değerlendirme yapabilirim, rock pop olmuştu. Gerçekten öz rock, hiç oryantal bir etkileşim olmadan esinlenme olmadan bangır bangır rock pop olmuştu. Sonradan müzik poplaşmaya başladı. Yükselen rock gruplarının şu an yaptığı, rock'ın geldiği hal kıvamı daha garip bir Türkçe Rrock var şu an ortada. Onun da sebebinin bu olduğunu düşünüyorum. Rrapin böyle olmasının sebebi insanlar artık bu sertliği ve ritmi ya da kendi kafasını sallatacak şeyi ve enerjisini atabilecek şeyi rapte buluyor. Şu anki hali enerjik değil ama rock'ın m. Modası geçmiş durumda. -Sözleri belli bir kalıp üzerine mi yazıyorsun? Belli bir tarzın ve tekniğin var mı? Yoksa o anda hissettiğin her neyse birden mi geliyor? Başka bir tarz da deneyebilir misin gün? Dediklerinin hepsinin toplamı gibi aslında. Mesela müzik yapacağız, ortada bir altyapı var veya hiçbir şey yok, o tamamen o anda yavaş yavaş oluşmaya başlıyor. Boş beat dediğimiz, sadece altyapı, söz yok bir şey yok, duyduğum zaman ilk aklıma gelen şey, bundan ne yapılabilir? Bu altyapıyla hislerim, sesim uyabilir mi? Ve dinlediğim anda şunu düşünmeye başlıyorum: nasıl bir nakarat çıkarabilirim? Bu şarkının hissiyatı ne? Bu hissiyata göre nasıl bir konu izleyebilirim, nasıl bir slogan bulabilirim? Şuna çok dikkat ederim, şarkıya girişin çok önemli olması lazım. Şarkıya girerken iyi girersen insanları tutarsın. Devamını dinlemeye devam ederler.

Böyle küçük şeylere dikkat ediyorum. Rapte zaten olan normal teknik olarak yapman gereken şeyler. “Flow” dediğimiz bir şey var; akıcılık, yani ritmin üstüne senin böldüğün ritimler. Bunların hepsinin teknik olarak mantıklı bir şekilde birleştiği bir şey yapmaya çalışıyorum ama bu süreçte tamamen içinden gelme, o ilham dedikleri de oluyor. Bir iki şeye dikkat edip sonra kalemi kendine bırakıyorum o şekilde gidiyor. Bazen oluyor koca bir yeri yazıyorum ama o son dörtlüğü yazamıyorum. Çünkü bakıyorum tam bu değil. Bbir hafta sonra yazıyorum belki. O an beş dakikada yazıyorum ama çok değişen bir şey. “Felaket”i mesela hiç altyapı olmadan enstrüman olmadan dolmuşta giderken mırıldanarak yaptım. Hayatımda ilk defa öyle bir beste yaptım ama ikinci verse’ünü aylar sonra Aankara’da yazdım mesela. İstanbul’daydım o zaman. Ankara’ya döner dönmez aldım gitarı elime, stüdyoyaya girdim çaldım gitarı, basları oluştı altyapı, üstüne kaydettim. Şarkı sana kendini yazdırıyor zaten.

Summaried Text

Son dönemlerin parlayan yıldızı ünlü rapçi Sercan İpekçioğlu, herkesin tanıdığı ismiyle "Ezhel", rap müziği ve kendisini anlattı.

Category

Kültür Sanat

Headline

Cayır Cayır Rap Yapıyorum

Appendix D. Five Examples of the Collected Data from the TR-News-Sum Dataset (Continues)

5th Example.
Main Text
İtalya ülkedeki bütün okulları ve üniversiteleri süresiz olarak kapatarak, 4 Mart Çarşamba gününden itibaren coronavirüsünün yayılmasını engellemek adına acil önlemler aldı. Avrupa'nın en çok etkilenen ülkesinde, ölüm sayıları giderek artıyor. Sivil Koruma Ajansının (Civil Protection Agency) yaptığı açıklamaya göre, İtalya'daki toplam ölüm sayısı 107 olurken, son 24 saat içerisinde virüs sebebiyle ölüm sayısı 28 kişi olarak belirlendi. Eğitim Bakanı Lucia Azzolina, “ Tüm ülkede, Perşembe gününden itibaren tüm okulların ve üniversitelerin en azından 15 Marta kadar kapalı olacak” dedi. Şu ana kadar sadece büyük risk altında olan ve fazlaca etkilenen Kuzey Bölgesinde okullar tatil edilmişti. Salgının 13 gün öncesinde ortaya çıkmasından bu yana vaka sayısı Salı günü 2.502'den 3.089'a yükseldi. Sivil Koruma Ajansı başkanı Angelo Borelli hastalığa yakalananların yaklaşık %3,5'inin öldüğünü söyledi. Hükümet günde yaklaşık 500 kişi artan vaka sayısından sonra bir kararname yürürlüğe koyarak önlemlerle alakalı durumun ne kadar ciddi seviyede olduğunu duyurdu. Başbakan Giuseppe Conte basına yaptığı açıklamada hastanelerdeki verimliliğe rağmen aşırı yoğunluk sebebiyle doktorların bunalmış olma riski altında olduğunu, yoğun bakım ünitelerinde sorunlar yaşadıklarını belirtti. Kararname, “İnsanların kalabalıklaşmasını ve en az bir metre yakınlık mesafelerini koruyamayacakları her türden olayın askıya alınması” maddesini içeriyor. Aynı zamanda sinemaya ve tiyatroların kapanması, insanlara sarılma veya el sıkışma gibi yakın temas gerektiren bütün fiziksel temaslardan kaçınmaları uyarıları kararname dahilinde. Okulların kapanması ise bazı öğrencilerde sevinç uyandırırken, ebeveynler durumla ilgili oldukça endişeliler. Roma'da bir ilkokul öğrencisinin velisi Massimiliano Del Ninno, “ Bu kararnameyi destekliyoruz, çünkü okullarda yayılabilecek olan herhangi bir salgın bizleri çok korkutuyor. Risk altında olmayan bir yaş grubu olsalar da, taşıyıcı olma ihtimalleri hala var” dedi.
Summaried Text
İtalya ülkedeki bütün okulları ve üniversiteleri süresiz olarak kapatarak, 4 Mart Çarşamba gününden itibaren coronavirüsünün yayılmasını engellemek adına acil önlemler aldı. Avrupa'nın en çok etkilenen ülkesinde, ölüm sayıları giderek artıyor.
Category
Sağlık
Headline
Coronavirüs Sebebiyle Ölenlerin Sayısı 107'ye Sıçradı, İtalya'da Okullar Tatil Edildi

Appendix E. Three Summarization Examples of Three Proposed Neural Network Models from the TR-News-Sum Dataset (Continues)

1st Example.
Main Text
Bu anlaşmada, ortalama küresel sıcaklık artışının, sanayileşme öncesi döneme göre 2 santigrat derecenin oldukça altında tutulması ve bunun en fazla 1,5 santigrat derece ile sınırlandırılması asıl amaç haline geldi. Gelişmiş ve gelişmekte olan ülkelerden periyodik olarak seçilen iki eş başkan önderliğinde anlaşmanın taslak hali oluşturulmuştu. 21. Taraflar Kongresi 2015 yılında Paris'te toplandığında son halini alan anlaşmanın, 2020 yılında yürürlüğe gireceği hedeflendi. Paris Anlaşmasını Kyoto Protokolünden ayıran en temel fark, ülkelerin sınıflandırılması konusu. Ülkeler zamanında oluşturulmuş ve güncelliğini yitirmiş gelişmişlik seviyesi sınıflandırılmasından arındırılmış, “Gelişmiş Ülkeler” ve “Gelişmekte Olan Ülkeler” olarak ikiye ayrılmıştır. Diğer bir önemli fark ise sera gazı emisyonlarının azaltılmasına yönelik. Kyoto Protokolünde ülke bazında sayısal bir hedef belirlenmişken, Paris Anlaşmasında ortalama küresel sıcaklık artışına yönelik hedefler vardır. Bu anlaşmada, ortalama küresel sıcaklık artışının, sanayileşme öncesi döneme göre 2 santigrat derecenin oldukça altında tutulması ve bunun en fazla 1,5 santigrat derece ile sınırlandırılması asıl amaç haline geldi. Yani temelde mesele, sadece gelişmiş ülkelerle çözüme ulaşmak değil, küresel iklim değişikliği riskine karşı tüm ülkeleri paydaş olarak ele alıp birlik olunması çağırısının yapılmasıydı. Meclisten Geçirerek Yürürlüğe Koymayan 10 Ülkeden Bir Tanesi Türkiye Paris İklim Anlaşması 7 Ekim 2019 tarihinde Rusya’da yürürlüğe girdi. New York’ta yapılan Birleşmiş Milletler Zirvesi sırasında anlaşmaya katılım sağlayacağı beyanında bulunan Rusya ile birlikte anlaşmaya katılım sağlayan ülke sayısı, Avrupa Birliği de sayılırsa 187’ye ulaştı. Rusya’nın da listeye eklenmesi ile birlikte dünya çapında, 197 ülke arasında anlaşmayı meclislerinden geçirerek yürürlüğe koymayan 10 ülke kaldı. Bunlardan bir tanesi ise Türkiye. 2016 Nisan ayında yapılan temsili imza gününe katılarak anlaşmayı imzalayan Türkiye, sonrasında niyet beyanlarında bulunurken politikalarını açıkladı. İklim şartlarını geri çevirmeye yönelik olmasa da olumsuz şartları durdurmaya yönelik alınacak olan aksiyon planını politikaları içerisinde açıklayan Türkiye, coğrafi konumu sebebiyle riskli bir noktada yer almaktadır. İzmir Yüksek Teknoloji Enstitüsü Çevre Mühendisliği bölümünde öğretim üyesi olan Profesör Doktor Aysun Sofuoğlu konu ile ilgili “Türkiye için yapılan risk analizinde Kuzey Afrika iklim kuşağının Türkiye’ye kayması sonucu kuraklık ve buna istinaden orman yangınlarında artış vs. gibi durumların özellikle Güneydoğu, Akdeniz ve Ege bölgesinde etkin olacağı risk olarak ifade edilmektedir.” diyerek bunların Çevre Bakanlığı raporlarında da yer aldığını belirtiyor. En büyük sorunlardan birinin kitlesel göçler olacağını, insanların yaşanılabilir yeni yerler arayacağını altını çizen Sofuoğlu, “Tarımda susuzluk nedeniyle rekolte düşmesi, ya da ürünü yetiştirmemenin ekonomik kayıplara neden olacağı açıktır.” diyerek konunun ekolojik sonuçlarının yanı sıra ekonomik sonuçlarına da değiniyor. Sofuoğlu, “Türkiye için konuşursak, ekolojik ve sosyoekonomik açıdan iklim değişikliği stresi karşısında da önlemler almıştır, belirlediği strateji ve politikaları mevcuttur” diyor ve ekliyor; “Türkiye’nin özellikle iklim değişikliğine uyum süreci altında belirlediği strateji ve politikaları mevcuttur. Ulusal İklim Değişikliği Uyum Stratejisi ve Eylem Planı ile Türkiye’de iklim değişikliğinden etkilenebilirlik alanlarını temel alarak: Su Kaynakları Yönetimi, Tarım ve Gıda Güvencesi; Ekosistem Hizmetleri, Biyolojik Çeşitlilik ve Ormancılık; Doğal Afet Risk Yönetimi; İnsan Sağlığı alanında çeşitli eylemler yer almaktadır. Örneğin atık suların arıtıldıktan sonra tarım ve sanayi sektöründe kullanılması teşvik edilmesi gibi...” “Başka Alternatifimiz Yok” Türkiye’nin belirlenen şartlara uyup uymayacağı konusunda, Ege Üniversitesi İktisadi İdari Bilimler Fakültesi İktisat Bölümü Öğretim Üyesi Doçent Doktor Meneviş Uzbay Pirili, “Dünya artık kapitalizm üzerinden ilerliyor. Dünyada önde gelen şirketlerin CEO’ları her sene Davos’ta toplanıp “World Economic Forum”u (Dünya Ekonomi Forumunu) gerçekleştiriyorlar. Bu insanlar için kapitalist toplumun kalbi diyebiliriz. Her sene bir risk raporu yayınlarlar ve bu risk raporunda hangi alanda yatırım yapılmalı,

riskler nedir gibi konular yer alır. Son raporlarında belirlenen en büyük risk çevresel riskler olarak belirlendi.” diyerek konunun ekolojik ve ekonomik kısmına dikkat çekti. “10 sene önceki raporda çevresel riskler en aşağılarda yer alıyordu. Bu kapitalizmin artık yeşil enerjiye, sürdürülebilir dünyaya, ekolojik sistemle uyumlu yatırımlara yöneldiğini gösterir. Dolayısıyla para ve teknoloji artık bu şekilde şekillendiği için Türkiye belirlenen şartları yerine getirmek zorunda. Başka alternatifimiz yok.” dedi.
Summarized Text from the Dataset Paris İklim Anlaşması dünya çapında 187 ülke tarafından imzalanarak 2020 yılında yürürlüğe girdi. Bu anlaşmada, ortalama küresel sıcaklık artışının, sanayileşme öncesi döneme göre 2 santigrat derecenin oldukça altında tutulması ve bunun en fazla 1,5 santigrat derece ile sınırlandırılması asıl amaç halindedir. Anlaşmayı imzalamayan 10 ülkeden biri Türkiye iken, anlaşmanın uluslararası güven çerçevesinde incelenmesiyle Türkiye sınıfta kalan ülkeler arasında yer almıştır.
Summarized Text by (i) Attention Based Seq2Seq Neural Network ConceptNet-Numberbatch sıcaklık <UNK> sanayileşme öncesi 2 <UNK> <UNK> altında <UNK> bunun 1 derece asıl <UNK>
Summarized Text by (ii) Pointer Generator Seq2Seq Neural Network ConceptNet-Numberbatch Ortalama küresel <UNK> <UNK> sıcaklık sanayileşme <UNK> öncesi göre 2 <UNK> oldukça altında bunun fazla 1 derece asıl amaç <UNK> <UNK>
Summarized Text by (iii) Reinforcement Learning Based Seq2Seq Neural Network ConceptNet-Numberbatch Ortalama sıcaklık <UNK> sanayileşme öncesi 2 <UNK> altında <UNK> 1 derece <UNK> asıl
Summarized Text by (i) Attention Based Seq2Seq Neural Network fastText küresel sıcaklık artışının sanayileşme öncesi döneme göre 2
Summarized Text by (ii) Pointer Generator Seq2Seq Neural Network fastText anlaşmada ortalama küresel sıcaklık artışının sanayileşme öncesi döneme göre 2 derece
Summarized Text by (iii) Reinforcement Learning Based Seq2Seq Neural Network fastText anlaşmada ortalama sıcaklık artışının sanayileşme öncesi döneme göre 2

Appendix E. Three Summarization Examples of Three Proposed Neural Network Models from the TR-News-Sum Dataset (Continues)

2nd Example.
Main Text
Güney Kore, Coronavirüs'ün (Covid-19) mücadelesinde Dünya Sağlık Örgütü tarafından örnek gösterildi. 2019 yılının Aralık ayında, Çin'de birdenbire sebebi açıklanamayan zatürre olguları ortaya çıkmaya başladı. Yapılan araştırmalar, bu zatürre vakalarının daha önceden tanımlanmamış yeni bir tip Coronavirüs olduğunu ortaya çıkardı. Coronavirüsler 21. Yüzyıl'da keşfedilmiş ve 1960'dan sonra insanlarda hafif solunum yolu hastalıklarına neden olduğu anlaşılmış, tüm dünyada yaygın bulunan virüslerdir. Güney Kore Tecrübeli Çin'den sonra Coronavirüs vakasının en çok görüldüğü ülkelerden biri olan Güney Kore; Dünya Sağlık Örgütü tarafından mücadelede de örnek gösterilmiştir. Güney Kore'de vaka sayısı yüksek olmasına rağmen yaşamını yitiren kişi sayısı oldukça düşük. Önlemler arttıkça, vaka sayılarında da oldukça düşüş görüldü ve bu şekilde virüse karşı ülkede etkili çalışmalar yürütüldü. Güney Kore, önceki yıllarda Mers hastalığından edindiği deneyimle Coronavirüs'e karşı başarılı bir mücadele verdi. Uzmanlar bu mücadelenin en önemli yardımcısının test olduğunu söylüyor. Güney Kore'de uzmanlar testlere yönelmişler ve test test test demişlerdir. Test kiti geliştirmek için de hızlı hareket etmişlerdir. Virüs şüphesi duyulan herkese test uygulamışlardır. Hafif semptom gösterenlerin kamuya açık alanlara çıkması engellenmiştir. Hızlı önlemlerle, hızlı bir gelişme göstermişlerdir.
Summarized Text from the Dataset
Virüsle Mücadelede Örnek: Güney Kore Güney Kore corona virüsün (Covid-19) mücadelesinde Dünya Sağlık Örgütü tarafından örnek gösterildi. Çin'den sonra corona virüs vakasının en çok görüldüğü ülkelerden biri olan Güney Kore; daha önceleri ülkelerinde görülen MERS hastalığından edindiği deneyimle, corona virüse karşı başarılı oldu. Ülkede vaka sayısı yüksek olmasına rağmen yaşamını yitiren kişi sayısı oldukça düşük.
Summarized Text by (i) Attention Based Seq2Seq Neural Network
ConceptNet-Numberbatch
örnek güney kore corona covid mücadele <UNK> <UNK>
Summarized Text by (ii) Pointer Generator Seq2Seq Neural Network
ConceptNet-Numberbatch
güney kore <UNK> <UNK> corona covid
Summarized Text by (iii) Reinforcement Learning Based Seq2Seq Neural Network
ConceptNet-Numberbatch
örnek güney kore <UNK> <UNK> güney kore corona covid
Summarized Text by (i) Attention Based Seq2Seq Neural Network
fastText
güney kore coronavirüs'ün örgütü kore dünya sağlık örgütü kore sonra
Summarized Text by (ii) Pointer Generator Seq2Seq Neural Network
fastText
güney kore coronavirüs'ün covid 19 mücadelesinde dünya sağlık örgütü tarafından
Summarized Text by (iii) Reinforcement Learning Based Seq2Seq Neural Network
fastText
covid 19 mücadelesinde dünya sağlık örgütü tarafından örnek gösterildi 2020

Appendix E. Three Summarization Examples of Three Proposed Neural Network Models from the TR-News-Sum Dataset (Continues)

3rd Example.
Main Text
Avrupa Uzay Ajansı'nın (ESA) ortaya koyduğu yeni bir çalışma ile Samanyolu Galaksisinin eğriliğinin zaman içerisinde değiştiğini keşfetti. Bu gelişmeler ile birlikte uzayı ve kendi evrenimizi tanımaya bir adım daha yaklaştık. Bu yolla yıllardır gizem olarak kalan ve evrenimizi tanımamıza büyük katkı sağlandı. Güneş sistemini de içinde barındıran Samanyolu Galaksisini tanımak ve haritalandırmak için Avrupa Uzay Ajansı'nın (ESA) başlattığı proje kapsamında Gai uydusunu kullanan uzmanlar sarmal galaksinin dönerken neden diğer galaksiler gibi davranmadığı sorusuna cevap getirdi. Bu gelişmeler sayesinde ileride yapılacak çalışmalar için yeni ufaklar ve çalışma alanları açıldı. Bilim insanları 1950'den beri tamamen düz olmadığını biliyorlar ancak bunun nedeni yapılan son çalışmalara kadar tam olarak bilinmiyordu. Ancak GAI uydusuyla yapılan son analizler sonucunda Samanyolu Galaksisindeki dalgalanmanın sebebi 11 milyar yıl önce galaksinin merkezine yakın bir yerde meydana gelen çarpışmadan kaynaklanıyor. Öte yandan Avrupa Uzay Ajansı'nın (ESA) daki uzmanlar yakındaki bir galaksinin çekim etkisinin eğrilğe neden olup olamadığı konusunda emin değiller. Çalışmanın başyazarı Turin Astrofizik Gözlemevi'nden Eloisa Poggio, "Verileri modellerimizle karşılaştırarak eğrilik hızını ölçtük bu yolla evrenimizi ve galaksi izi tanımaya bir adım daha yaklaştık ve bu gelişmeler oldukça heyecan verici" dedi. Ayrıca Uluslar Arası Uzay Ajansından (ESA) Poggio, elde edilen hıza olarak eğriliğin Samanyolu merkezinin etrafında dönüşünü 600 ila 700 milyon yılda tamamlayacağını söylerken elde edilen verilerin beklenenden çok daha fazla olduğunun ve bu veriler ışığında pek çok başka çalışmanın ve projenin de başlatılacağını da altını çizdi.
Summarized Text From the Dataset
Avrupa Uzay Ajansı'nın (ESA) ortaya koyduğu yeni bir çalışma ile Samanyolu Galaksisinin eğriliğinin zaman içerisinde değiştiğini keşfetti. Bu gelişmeler ile birlikte uzayı ve kendi evrenimizi tanımaya bir adım daha yaklaştık. Bu yolla yıllardır gizem olarak kalan ve evrenimizi tanımamıza büyük katkı sağlandı.
Summarized Text by (i) Attention Based Seq2Seq Neural Network
ConceptNet-Numberbatch
avrupa uzay ortaya yeni bir çalışma ile keşfetti <UNK>
Summarized Text by (ii) Pointer Generator Seq2Seq Neural Network
ConceptNet-Numberbatch
avrupa uzay ortaya <UNK> çalışma keşfetti
Summarized Text by (iii) Reinforcement Learning Based Seq2Seq Neural Network
ConceptNet-Numberbatch
avrupa <UNK> <UNK> uzay <UNK> <UNK> sağlandı <UNK>
Summarized Text by (i) Attention Based Seq2Seq Neural Network
fastText
avrupa uzay ajansı'nın ajansı'nın ajansı'nın ajansı'nın esa ajansı'nın esa başlattığı
Summarized Text by (ii) Pointer Generator Seq2Seq Neural Network
fastText
avrupa uzay ajansı'nın esa ortaya koyduğu yeni bir çalışma biçimi
Summarized Text by (iii) Reinforcement Learning Based Seq2Seq Neural Network
fastText
avrupa uzay ajansı'nın esa ortaya koyduğu yeni bir yöntem geliştirdiler