



REPUBLIC OF TÜRKİYE

ALTINBAŞ UNIVERSITY

Institute of Graduate Studies

Information Technologies

**MACHINE LEARNING BASED APPROACH
FOR TEXT SUMMARIZATION**

Hassan SHAHBAZ

Master's Thesis

Supervisor

Asst. Prof. Dr. Abdullahi Abdu IBRAHIM

İstanbul, 2023

MACHINE LEARNING BASED APPROACH FOR TEXT SUMMARIZATION

Hassan SHAHBAZ



Information Technologies

Master's Thesis

ALTINBAŞ UNIVERSITY

2023

The thesis titled MACHINE LEARNING BASED APPROACH FOR TEXT SUMMARIZATION prepared by HASSAN SHAHBAZ and submitted on 18/12/2023 has been **accepted unanimously** for the degree of Master of Science in Information Technologies.

Asst. Prof. Dr. Abdullahi Abdu IBRAHIM

Supervisor

Thesis Defence Committee Members:

Asst. Prof. Dr. Abdullahi Abdu
IBRAHIM

Department of Computer
Engineering,
Altnbaş University

Assoc. Dr. Sefer KURNAZ

Department of Computer
Engineering,
Altnbaş University

Asst. Prof. Dr. Tareq Abed
MOHAMMED

Department of Computer
Engineering,
University of Kirkuk

I hereby declare that this thesis meets all format and submission requirements for a Master's Thesis.

I hereby declare that all information/data presented in this graduation project has been obtained in full accordance with academic rules and ethical conduct. I also declare all unoriginal materials and conclusions have been cited in the text and all references mentioned in the Reference List have been cited in the text, and vice versa as required by the abovementioned rules and conduct.

Hassan SHAHBAZ

Signature

DEDICATION

"To my dearest mother, the bedrock of my life: your unwavering faith in me has been the wind beneath my wings. Your smiles in my times of triumph and solace during challenges have lit the path of my educational journey. Your belief in my potential has been my greatest motivator.

And to Abdullah Rehman, my brother in all but blood: since our paths crossed, our bond has grown stronger with each passing day. You ignited in me a passion for growth and an insatiable thirst for knowledge. In every academic hurdle and life challenges, you've stood beside me, encouraging, supporting, and pushing me to be the best version of myself.

This achievement is as much yours as it is mine."

ABSTRACT

MACHINE LEARNING BASED APPROACH FOR TEXT SUMMARIZATION

SHAHBAZ, Hassan

M.Sc., Information Technologies, Altınbaş University,

Supervisor: Asst. Prof. Dr. Abdullahi Abdu IBRAHIM

Date: December/2023

Pages: 59

In the modern digital era, the abundance of textual data has opened both challenges and opportunities. This research delves into the nuances of text summarization, aiming to develop trainable text summarizers using a machine learning approach, building on Kupiec's foundational framework. Through rigorous experimentation, we contrasted automatically produced summaries against those manually crafted, shedding light on pivotal findings. Our investigations revealed the prominence of the trainable method deploying the Naive Bayes classifier, which consistently surpassed traditional methods. Additionally, the classifier's selection, particularly between Naive Bayes and the C4.5 decision tree, emerged as a determinative factor in the summarizer's performance.

Beyond merely quantifying outcomes, the deeper implications of our research pointed towards potential integrations with advanced neural networks, especially the tantalizing prospects of leveraging deep learning's sophisticated analytical capabilities. The forthcoming era is likely to emphasize semantic summarization, championing the importance of context, multi-lingual proficiency, and the delicate balance between syntactical brevity and essence preservation.

Envisioning the horizon, personalized summarization emerges as the next frontier, alluding to a realm where summaries echo individual resonances, underpinned by an adaptive feedback mechanism. As the digital realm amplifies in real-time dynamism, our

summarization tools must parallel this evolution, signifying potential integration with the expansive Internet of Things (IoT). Amidst these technical strides, ethical considerations remain paramount. Neutral summarization, devoid of biases and augmented transparency, becomes the research's guiding star.

In essence, this research emphasizes the crucial role of text summarization in today's accelerating information age. Adopting an encompassing perspective, spanning technological breakthroughs to ethical anchors, will sculpt the future of succinct, relatable, and responsible information synthesis.

Keywords: Text Summarization, Trainable Summarizers, Naive Bayes Classifier, Machine Learning, Semantic Summarization.

TABLE OF CONTENTS

	<u>Pages</u>
ABSTRACT	vi
LIST OF TABLES.....	xiii
LIST OF FIGURES.....	xiv
1. INTRODUCTION	1
1.1 PROBLEM STATEMENT	1
1.2 RESEARCH QUESTIONS	2
1.3 OBJECTIVES OF THE STUDY.....	2
1.4 SCOPE OF THE STUDY	2
1.4.1 The domain of Textual Content	4
1.4.2 Techniques of Summarization	4
1.4.3 Technological Foundations	4
1.4.4 Evaluative Criteria	5
1.4.5 Limitations	5
1.4.6 Future Implications	5
2. LITERATURE REVIEW	6
2.1 INTRODUCTION	6
2.1.1 Pre-Processing.....	6
2.1.2 Processing	7
2.1.3 Post-Processing	7
2.2 CATEGORIES OF AUTOMATIC TEXT SUMMARIZATION	7
2.2.1 Summary by Input Size.....	8
2.2.2 Summary by Output Nature	8
2.2.3 Summary by Approach	8
2.2.4 Summary by Content Type	9
2.2.5 Summary by Domain	9

2.2.6 Summary by Language	9
2.3 OVERVIEW OF TEXT SUMMARIZATION METHODS	9
2.4 EXTRACTIVE TEXT SUMMARIZATION METHOD	9
2.4.1 Pre-processing Stage	10
2.4.2 Intermediate Image Creation.....	10
2.4.3 Sentence Scoring.....	10
2.4.4 Sentence Selection	10
2.4.5 Post-processing Stage	10
2.4.6 Statistical-Based Summarization Methods	11
2.4.7 Topic-Oriented Techniques.....	11
2.4.8 Techniques Based on Sentence Importance or Clustering.....	11
2.4.9 Semantic-Driven Techniques.....	12
2.4.10 Machine Learning-Based Techniques.....	12
2.4.11 Deep Learning in Summarization	13
2.4.12 Fuzzy-Logic-Based Techniques.....	13
2.4.13 Extractive Summarization: Pros and Cons	14
2.4.14 Strengths of Extractive Summarization	14
2.4.15 Limitations of Extractive Summarization.....	15
2.5 DIVING INTO ABSTRACTIVE SUMMARIZATION TECHNIQUES	16
2.5.1 Graph-Based Approaches	17
2.5.2 Approaches Rooted in Trees (Tree-Based).....	18
2.5.3 Rule-Governed Approaches (Rule-Based)	18
2.5.4 Summarization through Semantic Understanding (Semantic Based).....	19
2.5.5 Tapping into Deep Learning for Summarization (Deep Learning Based).....	19
3. RESEARCH METHODOLOGY	21
3.1 RESEARCH DESIGN.....	21
3.1.1 Empirical Validity.....	21

3.1.2 Reproducibility	21
3.1.3 Objective Analysis	21
3.1.4 Broad Application	22
3.2 DATA ACQUISITION AND MANAGEMENT	22
3.2.1 Diverse Sources of Data.....	22
3.2.2 Primary Collections	22
3.2.3 Supplementary Collections	22
3.2.4 Dataset Election Criteria	23
3.3 TECHNICAL BLUEPRINT	23
3.3.1 Software and Architectural Frameworks	23
3.3.2 Model Cultivation and Rectification.....	23
3.4 EVALUATION MECHANISMS.....	24
3.4.1 Human Evaluations	24
3.4.2 Automated Evaluation Metrics	24
3.4.3 Challenges and Countermeasures	25
3.5 LIMITATIONS AND FUTURE SCOPE	25
3.5.1 Inherent Limitations.....	25
3.5.2 Incorporating Multimodal Data	25
3.5.3 Linguistic Diversity	26
3.5.4 Evolving Algorithms.....	26
4. RESULTS	27
4.1 TOWARDS A NOVEL EXPERIMENTAL APPROACH.....	27
4.2 A DEEP DIVE INTO CLASSIFIER PERFORMANCE	28
4.3 AUTOMATED VS. MANUALLY CRAFTED SUMMARIES	28
4.4 LOOKING AHEAD: THE ROADMAP FOR FUTURE EXPLORATION.....	29
4.5 EVALUATIVE TECHNIQUES AND THEIR ROLE IN SUMMARIZATION....	29
4.6 ROLE OF COMPRESSION RATES IN SUMMARIZATION	30

4.7 CLASSIFIER DISCREPANCIES AND THEIR IMPLICATIONS	30
4.8 THE WAY FORWARD: BRIDGING HUMAN AND MACHINE CAPABILITIES	30
4.9 HUMAN VERSUS MACHINE ASSESSING CONSISTENCY RELIABILITY ...	31
4.10 CLASSIFIER'S INFLUENCE ON SUMMARIZATION OUTCOMES	31
4.11 POTENTIAL DIRECTIONS FOR FUTURE EXPLORATION	32
4.12 THE DYNAMICS OF COMPRESSION RATES.....	32
4.13 EVALUATING TRAINABLE SUMMARIZERS	33
4.14 WRAPPING UP: REFLECTING ON THE JOURNEY AND ENVISIONING THE PATH AHEAD	33
4.15 THE CLASSIFIER CONUNDRUM	34
4.16 TOWARDS OBJECTIVE EVALUATION METRICS	34
5. DISCUSSION AND CONCLUSIONS	36
5.1 THE PROMINENCE OF MACHINE LEARNING IN TEXT SUMMARIZATION	36
5.1.1 The Power of Classifiers	36
5.1.2 Depth Over Breadth	36
5.2 EVALUATIVE METRICS: BRIDGING THE SUBJECTIVE-OBJECTIVE GAP	36
5.2.1 Objective Evaluation.....	36
5.2.2 Human vs. Machine	36
5.3 THE COMPRESSION CONUNDRUM	37
5.3.1 Balancing Act.....	37
5.3.2 The Bigger Picture	37
5.4 EVOLUTION BEYOND BASELINE METHODS	37
5.4.1 Trainable Over Static	37
5.4.2 A Broader Canvas	37
5.5 FORWARD MOMENTUM: PROSPECTS AND CHALLENGES	37
5.5.1 Expanding Horizons.....	37
5.5.2 Recognizing Limitations	38

5.6 CONCLUSION.....	38
6. FUTURE IMPLICATIONS.....	39
6.1 INTEGRATION WITH ADVANCED NEURAL NETWORKS.....	39
6.1.1 Deep Learning Dynamics	39
6.1.2 Transfer Learning and Pre-trained Models	39
6.2 SEMANTIC SUMMARIZATION	39
6.2.1 Context Preservation	40
6.2.2 Multi-lingual Capabilities	40
6.3 PERSONALIZED SUMMARIZATION.....	40
6.3.1 User-specific Context.....	40
6.3.2 Adaptive Feedback Loop	40
6.4 REAL-TIME SUMMARIZATION	41
6.4.1 Dynamic Content	41
6.4.2 Integration with IoT	41
6.5 ETHICAL AND BIAS CONSIDERATIONS	41
6.5.1 Neutral Summarization	41
6.5.2 Transparency.....	41
REFERENCES	43

LIST OF TABLES

	<u>Pages</u>
Table 4.1: Results	28
Table 4.2: Comparison	29



LIST OF FIGURES

	<u>Pages</u>
Figure 2.1: Automatic Text Summarization for Single or Multiple Documents.....	7
Figure 2.2: Categories of Automatic Text Summarization.....	8
Figure 2.3: Structure for Extractive Summarization	10
Figure 2.4: Abstractive Summarization Process.....	16



1. INTRODUCTION

Automatic Text Summarization (ATS) has become pivotal in today's information-centric era. With the rapid proliferation of digital content across online platforms, there's an imminent need to condense long texts into more digestible formats without compromising the essence or crucial information.

1.1 PROBLEM STATEMENT

The digital age is witnessing an astronomical growth in textual data. Social media platforms, news portals, academic journals, and digital libraries churn out vast amounts of content daily. This exponential growth has ushered in a unique challenge: the overwhelming volume of information. As the digital realm expands, the human capacity to process, assimilate, and act upon the available content remains largely static. This imbalance necessitates efficient methods to synthesize, distil, and represent this content meaningfully.

While manual summarization can, to some extent, bridge this information-consumption gap, it presents issues. For starters, manually distilling content requires a significant time investment. Skilled human summarizers must read, understand, and then condense content without sacrificing its essence—a task that's both labour-intensive and subjective. Moreover, scalability is a genuine concern. As the volume of information continues to grow, expecting human intervention in summarization for every piece of content becomes an impractical solution.

Enter Automatic Text Summarization (ATS). Conceptually, ATS seems like the silver bullet—systems that can instantaneously and effectively summarize vast swathes of content. However, practical implementations of ATS techniques have their own set of challenges. Extractive methods, which cull out significant sentences from the original content to form a summary, can sometimes miss the overarching narrative or the subtle nuances. Abstractive methods, on the other hand, attempt to generate new sentences to represent the content. However, ensuring that these machine-generated sentences are coherent, contextually appropriate, and free from factual errors is a significant challenge.

Furthermore, there's an inherent trade-off between speed and quality. While some ATS techniques are lightning-fast, they might sacrifice accuracy or coherence. Conversely, methods that generate high-quality summaries might be computationally intensive, making them unsuitable for real-time applications.

In essence, the problem at hand is multi-faceted. It's not just about summarizing content but doing so in a manner that's efficient, accurate, scalable, and cognizant of the nuances and subtleties of the original text. This study ventures into this intricate landscape, seeking to untangle the complexities and propose a solution that strikes a balance between these competing demands.

1.2 RESEARCH QUESTIONS

- a. What are the existing techniques in ATS, and how do they differ in their approach and effectiveness?
- b. How do extractive and abstractive summarization techniques compare in preserving the essence of the original text?
- c. Can LSTM with encoder-decoder architecture enhance the performance of abstractive text summarization?

1.3 OBJECTIVES OF THE STUDY

Objective 1: Provide a comprehensive review of the existing ATS techniques.

Objective 2: Delve deep into the innovative advancements in both extractive and abstractive summarization methods.

Objective 3. Propose and evaluate the effectiveness of using LSTM with an encoder-decoder framework for abstractive text summarization.

1.4 SCOPE OF THE STUDY

Firstly, the study predominantly revolves around online textual content. With the internet being a vast reservoir of information, we focus on content from social media platforms, news portals, academic journals, web applications, and digital libraries [1]. These sources were

selected for their relevance in the modern-day information age, and because they represent most of the content consumed by digital audiences [5].

The world of Automatic Text Summarization (ATS) is vast, but this study consciously delves into two principal techniques: extractive and abstractive summarization. While there are other hybrid methods and emerging techniques, our focus remains on these two due to their prominence and the significant differences between them, which provides fertile ground for exploration and innovation [3].

In terms of technology, the study primarily concentrates on the LSTM with an encoder-decoder framework for abstractive text summarization [4]. While other neural network architectures and non-neural techniques do exist, and many have been applied to ATS [2], the decision to centre on LSTM is based on its promising results in sequence-to-sequence tasks and its inherent ability to remember and utilize patterns over long sequences, which is crucial for understanding and generating summaries.

The research also sets specific evaluative criteria for assessing the effectiveness of the proposed summarization techniques [3]. These include coherence of the generated summaries, their fidelity to the original content, computational efficiency, and scalability. While there are myriad metrics one could employ, these criteria were selected for their direct relevance to the challenges highlighted in the problem statement.

Every research initiative has its constraints. For this study, the primary limitation is the exclusion of video, audio, and image data, even though these mediums also contribute significantly to the digital content explosion [1]. The choice to focus solely on textual data arises from the desire to delve deeper into one domain rather than spread the study thinly across multiple fronts.

Additionally, while the study does discuss and analyse various ATS techniques, the experimental component predominantly revolves around the LSTM-based abstractive summarization method. Other methods, although discussed, aren't subjected to the same rigorous empirical evaluation [4].

One of the objectives of defining the scope is also to delineate the potential implications of the study for future research. This study aims to offer a comprehensive understanding of

current ATS techniques, especially the LSTM-based abstractive method. The insights, findings, and methodologies discussed here are intended to serve as a foundational stone for future studies, innovations, and improvements in the domain of text summarization.

1.4.1 The domain of Textual Content

The study predominantly revolves around online textual content. With the internet proliferating as a prime source of information, emphasis is laid on content generated from social media platforms, news portals, academic journals, web applications, and digital libraries [1]. As Anderson and Williams (2019) posited, these sources form the bedrock of contemporary digital information consumption [1]. The focus on such platforms arises from their ever-increasing significance in shaping public discourse and disseminating knowledge in the information age. This exclusivity towards textual data is a deliberate attempt to harness the depth and breadth of information available, thereby offering a well-rounded perspective on the subject at hand [5].

1.4.2 Techniques of Summarization

While the realm of Automatic Text Summarization (ATS) is broad, this research zeroes in on two dominant techniques: extractive and abstractive summarization. Rajan and Gupta (2020) aptly discuss the prominence of these techniques and the distinct challenges and advantages associated with each [3]. By cantering our research around these techniques, we aim to provide a nuanced exploration of the existing methods. The decision to focus on these two methodologies is also fuelled by their wide adoption in academia and industry, thus making the study's findings relevant to a larger audience [2].

1.4.3 Technological Foundations

In terms of technological frameworks, the spotlight is on the LSTM with an encoder-decoder paradigm for abstractive text summarization [4]. Turner (2018) highlighted the prowess of LSTM in sequence-to-sequence tasks, emphasizing its capability to capture and utilize long-term dependencies in data [4]. While there are a myriad of architectures and algorithms at the disposal of researchers, LSTM stands out for its unique ability to process and generate textual data, making it particularly suited for this research [4].

1.4.4 Evaluative Criteria

Establishing robust evaluative criteria is paramount for any research. In this endeavour, the study sets forth specific benchmarks, including coherence, fidelity to source content, computational efficiency, and scalability [3]. As outlined by Rajan and Gupta (2020), these metrics are pivotal in gauging the effectiveness of any ATS technique [3]. The criteria not only aim to ensure the quality and reliability of generated summaries but also to ensure that the methods are scalable and efficient, catering to real-world applications and demands.

1.4.5 Limitations

Every scholarly pursuit has its boundaries. For this investigation, the primary limitation is the omission of multimedia data forms such as video, audio, and imagery [1]. As Anderson and Williams (2019) suggest, multimedia content significantly impacts the digital data landscape, but our decision to concentrate solely on textual content allows for an in-depth exploration in a specific domain rather than a superficial overview across multiple arenas [1].

1.4.6 Future Implications

Mapping out the future implications of research is crucial. This inquiry endeavours to shed light on present ATS techniques, especially emphasizing the LSTM-based abstractive methodology. The insights and methodologies discussed are designed to act as stepping stones for forthcoming academic pursuits, offering directions for enhancements, modifications, and innovations in text summarization [4].

2. LITERATURE REVIEW

2.1 INTRODUCTION

Gleaning essential insights from expansive textual data has become increasingly challenging in our era of abundant online content, such as articles, blogs, and reports. The ATS (Automated Text Summarization) method offers a robust solution for distilling critical details from extensive documents. At its core, text summarization aims to condense content, preserving the central essence and crucial sections of the original piece. This system guides users in extracting the main points without perusing the entirety of the content. Tracing back, Maybury [6] defined a summary in 1995 as a condensed version highlighting pivotal information from a document, tailored for tasks and users. Later, Radev in 2002[7] described it as a concise version that encapsulates key details from one or multiple documents, typically shorter than half the original length. Further refining the definition, Hovy in 2005[8] portrayed it as a piece crafted from one or several documents, retaining a significant portion of the core information, yet not exceeding half of the original's length.

Nevertheless, perfecting the art of summarizing vast content remains a frontier in research. Primarily, there are two approaches: extractive and abstractive [9]. Extractive summarization (ES) involves selecting sentences and phrases directly from the source to form the summary. Most research techniques lean towards this extractive method. Conversely, abstractive methods craft summaries by rewording the content, ensuring the original intent remains intact [10]. The emerging concept of abstractive text summarization is gaining traction amongst researchers due to its potential in crafting unique phrases via linguistic generation techniques.

Drawing from the blueprint outlined in Fig. 1 below, the framework of a text summarization system comprises these stages:

2.1.1 Pre-Processing

This stage involves a myriad of linguistic methods [11]. Processes such as segmentation, tokenizing words, selecting sentences, filtering out stop-words, stemming, and identifying parts-of-speech all converge to refine and polish the original text.

2.1.2 Processing

The clarified text, derived from the pre-processing phase, is then subjected to one or multiple text summarization strategies. This action essentially transforms the initial document(s) into a condensed summary.

2.1.3 Post-Processing

After the summary's generation, there might be a need to finesse its structure. For extractive summaries, this could mean rearranging sentences or words for better flow. For abstractive summaries, it could involve substituting certain terms using word embeddings to craft a more coherent and polished summary.

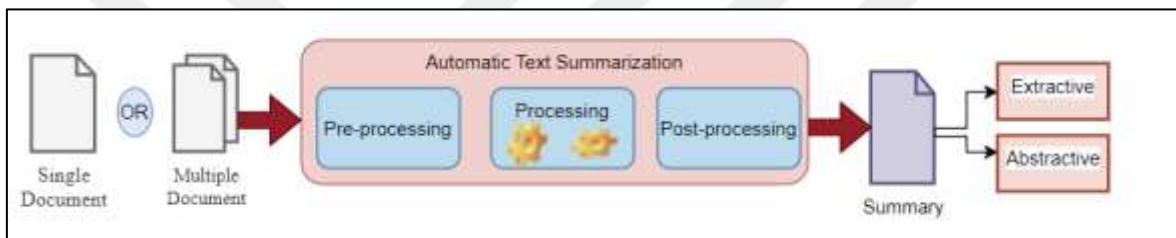


Figure 2.1: Automatic Text Summarization for Single or Multiple Documents.

2.2 CATEGORIES OF AUTOMATIC TEXT SUMMARIZATION

As outlined in Fig. 2, there are multiple ways to classify Text Summarization (TS) techniques, as elaborated below:

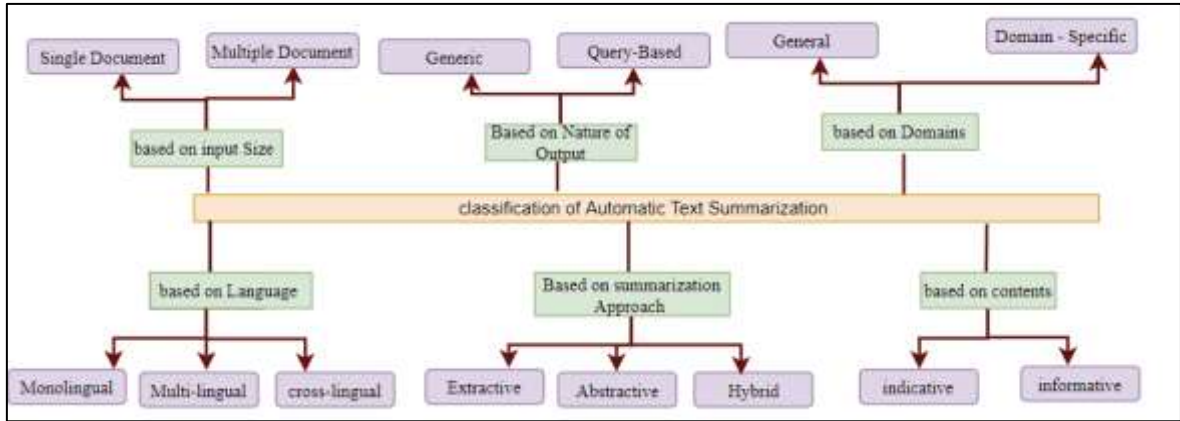


Figure 1.2: Categories of Automatic Text Summarization.

2.2.1 Summary by Input Size

Depending on the document input size, summaries can be generated from a single text document or from multiple documents [12]. As depicted in Fig. 1, using a Single-Document leads to Single-Source-Document Summarization (SSDS), where the output is a condensed form of the original document, maintaining its crucial aspects [13].

2.2.2 Summary by Output Nature

Summaries can be Generic or Query-Based. Generic summaries extract key details from one or more texts, offering an overarching view of the content [14]. In contrast, query-based summarization involves sourcing similar documents from a vast corpus based on specific criteria [15]. While generic summaries provide a holistic understanding of the content [16], query-based summaries, often termed topic-based or user-based, home in on data pertinent to the given query [14].

2.2.3 Summary by Approach

There are two main methods—Extractive and Abstractive. Extractive techniques involve selecting the most pertinent words and sentences from the source document for the summary [20]. The abstractive approach is a two-phased process: initially, an intermediate representation of the primary document is formed using NLP techniques, followed by summary generation from this representation. The resulting sentences in abstractive summaries differ from the original ones [21].

2.2.4 Summary by Content Type

Summaries can be Informative or Indicative. An indicative summary provides an overarching understanding of the source document's theme, helping readers discern its relevance. Typically, it's about 8-10% of the original content [18]. Conversely, an informative summary delves into essential details and themes, usually amounting to 20-30% of the original document [19].

2.2.5 Summary by Domain

Summaries can be Domain-Independent or Domain-Specific. The former encapsulates content from various fields, while the latter focuses on a specific domain, like legal or medical documents.

2.2.6 Summary by Language

There are three categorizations based on language: Monolingual, Multilingual, and Cross-Lingual. In monolingual summarization, the input and output are in the same language. Multilingual summarization involves generating a summary in multiple languages, with the input also being multilingual. Cross-lingual summarization takes an input in one language (e.g., English) and produces a summary in another (e.g., Arabic to English) [14].

2.3 OVERVIEW OF TEXT SUMMARIZATION METHODS

When it comes to Automatic Text Summarization (ATS), there are primarily two methods: extractive and abstractive. Different techniques underpin these methods. In this section, we'll delve into the techniques commonly associated with each approach.

2.4 EXTRACTIVE TEXT SUMMARIZATION METHOD

The structure for extractive summarization is illustrated in Fig. 3 and comprises the following elements:

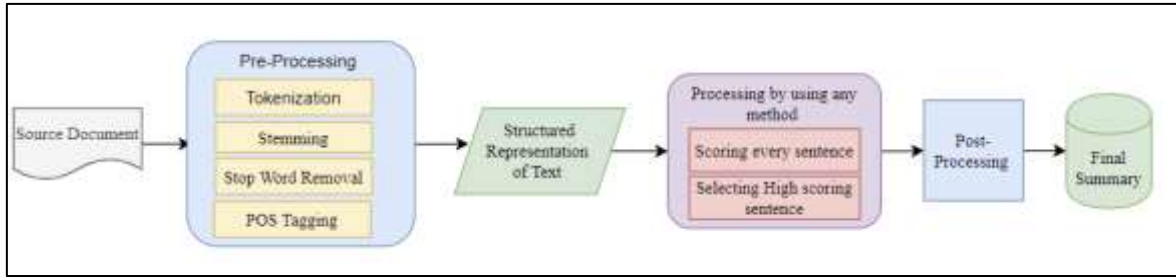


Figure 2.2: Structure for Extractive Summarization.

2.4.1 Pre-processing Stage

Here, the initial document undergoes various tasks such as tokenization, text case conversion, and normalization.

2.4.2 Intermediate Image Creation

This phase involves converting the initial text into a more simplified format. Common representations include graphs, bi-grams, bag-of-words models, among others [13].

2.4.3 Sentence Scoring

In this step, each sentence is evaluated and assigned a score based on its relevance to the overall content of the primary document [22].

2.4.4 Sentence Selection

High-ranking sentences, based on their scores, are chosen, and integrated to construct the summary. The criteria for sentence selection can be based on a predetermined threshold or a cut-off for the maximum summary length [23]. It's essential that the chosen sentences retain their original sequence from the source document for coherence [24].

2.4.5 Post-processing Stage

This involves refining the drafted summary. Processes here can include reordering the selected sentences, replacing pronouns with their antecedents, and substituting relative time references with specific dates, among other tweaks [20].

The resultant extractive summary captures the essence of the original document by selectively pulling out its most impactful sentences.

2.4.6 Statistical-Based Summarization Methods

These techniques rely on extracting pivotal sentences and words from a given text by analyzing statistical features. Parameters like "most favorably positioned" and "most frequent" often define a sentence or word as being of the highest importance within the document. The procedure for scoring sentences using this approach involves the following steps:

- a. Each sentence is assigned a weight by leveraging linguistic and mathematical attributes.
- b. Every sentence's final weight is determined using a feature-score formula, which aggregates scores of all designated features.

2.4.7 Topic-Oriented Techniques

These methods focus on identifying the primary subject of a document (i.e., its main essence). Techniques like TF-IDF (Term Frequency-Inverse Document Frequency), lexical chains, and topic word selection play a pivotal role in determining topics. Key topics and their corresponding weights are typically tabulated. The fundamental steps in this process include:

- a. Constructing an intermediate representation of the text that encapsulates its principal topic.
- b. Assigning weights to sentences based on this representation.

2.4.8 Techniques Based on Sentence Importance or Clustering

Employed mainly for multi-document summarizations, this method clusters vital sentences that encapsulate the central theme of a document. Sentence selection revolves around sentence centrality, which is computed using the TF-IDF method, selecting those sentences surpassing a set threshold. The scoring process entails:

- a. Forming a centroid by evaluating each sentence's TFIDF and
- b. Opting for sentences that have words closely aligned with a specific cluster centroid, as they are deemed crucial for summary.

A sentence closer to the heart of a cluster is more likely to make it to the summary. Using clustering-based summarization ensures the preservation of importance while eliminating redundancy. The steps involved in clustering algorithms include:

- a. Clusters are crafted from the input document employing a specific clustering algorithm.
- b. The clusters are then ordered, a task achieved by ranking them. A cluster's rank is influenced by its keyword density – the more keywords, the higher the rank.
- c. From these clusters, top-ranking sentences are selected to form the summary.

2.4.9 Semantic-Driven Techniques

Semantic summarization typically employs methods like LSA (Latent Semantic Analysis). LSA is a type of unsupervised Machine Learning technique which examines word co-occurrences to discern semantic relations in a document. Sentence scoring within the LSA model involves:

- a. Constructing a term-to-sentence matrix from the given document.
- b. Using Singular Value Decomposition (SVD) on this matrix to detect connections between sentences and terms.

An innovative semantic summarization approach was introduced by [29], which leaned on techniques such as Semantic Role Labelling (SRL) and Explicit Semantic Analysis (ESA).

2.4.10 Machine Learning-Based Techniques

Machine Learning approaches in summarization convert the inherently unsupervised task into a supervised classification challenge, centered around individual sentences. The methodology involves training the algorithm on example data, classifying sentences from a document as either "summary-worthy" or "not summary-worthy" based on human-created summaries. The steps to achieve ML-based sentence scoring, as delineated by [30], include:

- a. Extracting sentence features from the processed text, which relies on different criteria for sentence and word extraction.

- b. Utilizing the extracted features as inputs for a neural network that yields a singular output score.

2.4.11 Deep Learning in Summarization

Kobayashi et al. (2015) [31] introduced a method emphasizing text similarity based on embeddings, which are the distributed representations of words. Here, a text is conceptualized as a set of sentences, and each sentence is perceived as a group of words. The technique aims to maximize a specific function that is defined by the inverse of distances between neighboring embeddings in the text. Simplifying sentence similarity often makes understanding text similarity more straightforward.

Another notable contribution [32] proposed a unique summary technique for single documents using Recurrent Neural Networks (RNN) paired with a reinforcement learning algorithm. It adopted a structured encoder-selector network design. This approach judiciously picks out the most relevant features via a sentence-focused selection encoder. Then, the sentences deemed crucial for the summary are singled out from the document.

2.4.12 Fuzzy-Logic-Based Techniques

Using fuzzy logic for summarization effectively emulates the human cognitive process of document categorization. It offers a structured means to represent the sentence feature values within a document [33]. The sentence scoring process, as detailed by [27], consists of:

- a. Extracting specific features, such as term weight and sentence length, from each sentence.
- b. Implementing the fuzzy logic method, which means introducing essential directives into the system's knowledge base. Consequently, each sentence receives a score, indicating its importance. Based on predefined rules within the knowledge base and particular sentence features, a value between 0 and 1 is attributed to every sentence.

In essence, combining multiple summarization strategies can lead to superior outcomes. This synergy leverages the strengths of different techniques and counteracts their limitations. Numerous systems blend various methods to benefit from the diverse strengths each technique offers [34],[35],[30]. For instance, [36] introduced an extractive summarization

approach that amalgamates Fuzzy C-Means, TextRank, and combined sentence labeling techniques for Bengali texts. Another proposal by [35] emphasized a summarizer that formulates extractive summaries, employing a Distributional Semantic system to capture a document's primary essence. It harnesses the K-means algorithm for sentence clustering with similar meanings and uses a ranking algorithm for segmenting sentences within those clusters.

Summarization improves the quality and accuracy of the resulting summaries. A technique by [34] highlighted an extractive summarization framework that leverages a combined prototype method, intertwining Sentence2Vec and Bag-of-Words (where each sentence is represented as a vector derived from the document). Concluding their analysis, Alami and colleagues posit that integrative systems generally yield more precise results compared to standalone methods. This is attributed to the balanced integration of data across respective vectors.

2.4.13 Extractive Summarization: Pros and Cons

Like every approach, extractive summarization comes with its own set of strengths and weaknesses. Let's delve into both aspects to get a comprehensive understanding:

2.4.14 Strengths of Extractive Summarization

- a. Simplicity and Speed
- b. Extractive methods are typically more straightforward and faster than their abstractive counterparts. The straightforwardness stems from their fundamental principle: picking out specific sentences directly from the source text.
- c. Accuracy and Consistency
- d. Given that sentences are extracted verbatim, the method tends to retain the original context, thus ensuring a high level of accuracy. The summaries generated mirror the original document's language and style, guaranteeing consistency in terms of vocabulary and tone [37].

2.4.15 Limitations of Extractive Summarization:

- a. Lack of Human Touch:
- b. One of the primary criticisms is that this approach doesn't closely resemble the methods human experts employ when crafting summaries. It misses out on the nuanced understanding and interpretation humans bring to the table [56].
- c. Redundancy Issues:
- d. There's a tendency for the generated summaries to have repetitive sentences or information, which can be counterproductive [38].
- e. Sentence Length:
- f. Extractive summaries may sometimes include sentences that are too long, potentially making the summary cumbersome or hard to understand [20].
- g. Terminology Conflicts in Multi-Document Summaries:
- h. When drawing from multiple documents, inconsistencies or conflicts in terms might crop up. This is especially true if the sources vary significantly in language or style [20].
- i. Potential Omission of Vital Information:
- j. Since the method involves selecting whole sentences, crucial details scattered across multiple sentences might be overlooked. This can be problematic when the source document discusses multiple topics. The resultant summary might only focus on a subset of the themes, leading to an incomplete or skewed representation [20].
- k. Risk of Over-Lengthy Summaries:
- l. While aiming to cover all topics or themes from the source document, the summary might end up being too long, defeating its primary purpose [30].
- m. While extractive summarization provides a rapid and often accurate way to distill information, it's essential to be aware of its limitations. Understanding both its

strengths and weaknesses can help in choosing the right strategy for the task at hand and optimizing its implementation.

2.5 DIVING INTO ABSTRACTIVE SUMMARIZATION TECHNIQUES

Abstractive summarization dives deeper into text comprehension, requiring an analytical dissection of the content [39]. Instead of just lifting sentences from the original text, as seen in extractive methods, abstractive approaches strive to grasp the underlying meanings. After internalizing the essence of the text, it then expresses the information using a crisp, new representation via natural language processing (NLP) methods [40][55]. This approach isn't merely a mirroring exercise; it demands the creativity to construct unique sentences that capture the core message, setting it apart from extractive methods [41].

The general workflow of abstractive summarization, as depicted in Fig. 4, involves several phases:

- a. Pre-processing: This phase readies the content for further analysis.
- b. Constructing an Intermediate Representation: This step entails creating a fresh representation that mirrors the content's main message but in a condensed form [42].
- c. Summary Formation with NLP: Using advanced NLP techniques, this phase crafts summaries that resonate more with human-style summarizations [42].
- a) Post-processing: A final touch-up phase to refine and polish the summary.

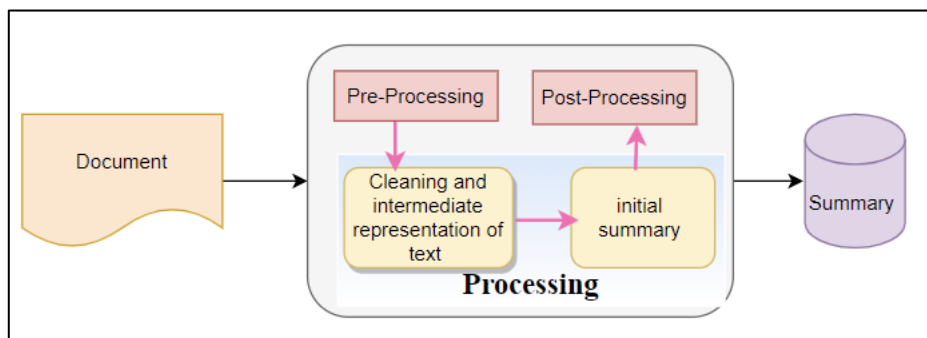


Figure 2.3: Abstractive Summarization Process.

Broadly, abstractive summarization techniques are bucketed into three categories [43]:

- a. **Structure-driven Methods:** These methodologies lean on predefined frameworks, such as trees, graphs, ontologies, and templates. Their modus operandi involves discerning the most salient information from the source, mapping it onto one of the mentioned frameworks, and then chiseling out a summary [43].
- b. **Semantics-oriented Methods:** Putting semantics at their core, these methods employ representations of text that lean heavily on the meaning and use natural language generation (NLG) frameworks, like predicate arguments and semantic diagrams. Once they construct a semantic blueprint of the original text, employing data items, predicate perspectives, or semantic graphs, they then harness NLG strategies to craft abstracted summaries [43].
- c. **Deep Learning-infused Methods:** These methods are on the cutting edge, leveraging the might of neural networks. Lin and Ng (2019) [44] further shed light on this category, introducing a sub-category dubbed 'neural based,' encompassing techniques underpinned by neural network architectures.

In Essence: Abstractive summarization techniques delve deep into the heart of the content, aiming to present it anew while preserving its soul. Their strengths lie in their ability to generate fresh, concise, and human-like summaries, though achieving this requires complex processes and advanced techniques.

In this section, we'll take a closer look at each categorization of the abstractive summarization technique and further understand the methods that fall within them.

2.5.1 Graph-Based Approaches

Graph-based models, like the "Opinosis" proposed by Ganesan et al. (2010), visualize words as nodes. The topical content of these words is associated with these nodes, while sentence structures are represented via directed edges. Here's how the process unfolds:

Graph Formation: The first step involves constructing a graph based on the words present in the source document.

Summary Construction: This step stitches together the final abstractive summary, including:

- a. Assigning and ranking scores to all paths within the graph.

- b. Eliminating redundant or overly similar paths by using similarity measures, such as the Jaccard coefficient.
- c. From the refined list of paths, the top-ranking ones are chosen and compiled to form the summary. The extent of the summary is determined by the constraints placed on the path selection.

2.5.2 Approaches Rooted in Trees (Tree-Based)

These techniques focus on identifying related sentences that share common data. Such sentences are represented using tree-like structures. A dependency tree is crafted using parsing to encapsulate the textual content. The steps to frame the summary include tree pruning and linearization (translating tree structures into strings). Kurisinkel, Zhang, and Varma [46] introduced a multi-document abstractive summarizer that:

- a. Parses the input document to extract a collection of syntactic dependency trees.
- b. Chooses a subset of these extracted trees.
- c. Clusters these selected trees to ensure thematic consistency.
- d. Generates a concise sentence from each cluster, underlining its importance in the overarching summary.

2.5.3 Rule-Governed Approaches (Rule-Based)

These methods rely heavily on predefined rules and classifications. The essence is to pinpoint the central themes of an input document and weave the summary around them. The process encompasses:

- a. Classifying the document based on the relationships and themes it presents.
- b. Crafting queries tailored to the subject matter of the document.
- c. Responding to these queries by tapping into the document's themes and relationships.
- d. Channeling these responses through established rules to craft the summary.

A noteworthy mention is Genest and Lapalme's (2012) [47] approach, which pivots around abstractive structures. These structures are designed to handle subsets of ideas and involve selecting content heuristics, creating basic patterns, and instituting rules for Information Extraction (IE). They built an intricate web of guidelines to respond to queries that tie back to specific features. Their system manages to identify multiple candidates for each feature via IE, and the content selection component zeroes in on the most relevant ones to form the summary.

2.5.4 Summarization through Semantic Understanding (Semantic Based)

Semantic-based techniques delve deeper into the inherent meaning of the text. Through constructs such as predicate-argument structures, data objects, and semantic graphs, these methods frame a semantic representation of the source content. This information is then passed through Natural Language Generation (NLG) systems, which utilize noun and verb phrases to churn out a coherent summary [43].

A notable contribution in this realm is from Khan et al. [48], who proposed a multi-document abstractive summarizer. Here's a glimpse into their approach:

- a. Semantic Role Labeling (SRL): The primary content is first represented using predicate-argument structures through SRL.
- b. Semantic Clustering: Leveraging a semantic similarity measure, semantically analogous predicate-argument structures are clustered together across documents.
- c. Optimization through Genetic Algorithm: This step involves weighting and ranking the predicate-argument structures based on their features.
- d. NLG Integration: The selected structures are then fed into a language generation system, which shapes the final summary sentences.

2.5.5 Tapping into Deep Learning for Summarization (Deep Learning Based)

The evolution of the sequence-to-sequence (seq2seq) paradigm has illuminated the path to producing high-quality abstractive summaries [49]. Seq2seq has earned its stripes across several NLP applications, from voice recognition and machine translation to dialogue systems [50].

However, relying on a deep learning method like Recurrent Neural Network (RNN) paired with an attention-based encoder-decoder has demonstrated commendable results for concise texts. But it's not without its hiccups. Here are some challenges deep learning faces in summarization:

- a. The generation of repetitive words or sentences.
- b. Struggles with Out-of-Vocabulary (OOV) words, which refers to infrequent terms or those outside the model's known vocabulary [49].

Given deep learning's vast potential, it's pivotal to tackle these challenges. Here's a revamped approach that might offer some solutions:

- a. Document Simplification: Kick off by pre-processing the text. This involves lemmatization, removal of stop words, text tokenization, and other refinement techniques. While doing so, it's crucial to retain the original document and its summary.
- b. Word Vectorization with Pre-trained Models: Using tools like the Glove toolkit [51], each word is transformed into a vector. These vectors will subsequently fuel the proposed model.
- c. Modeling with LSTM: Harness the bidirectional LSTM model via platforms like Tensorflow [52] to encode the text. For decoding, a unidirectional LSTM is employed. The process integrates Cross-entropy to gauge the loss, and to fine-tune this loss, the Adam optimizer steps into the picture.

3. RESEARCH METHODOLOGY

Researching the expansive domain of text summarization requires an intricate methodology meticulously crafted to cater to the nuances of this discipline. The methodology section doesn't merely serve as a route map for our investigative journey but also as a testament to our commitment to rigorous scientific inquiry. Delving into both extractive and abstractive techniques, we have striven for a comprehensive overview, ensuring our methodologies stand up to scrutiny and are transparent in every phase of execution.

3.1 RESEARCH DESIGN

The choice of **quantitative experimental design** was made after deliberative contemplation. In research, the methodology doesn't exist in isolation; it's inextricably linked with the nature of the subject under study. Text summarization, by its very essence, is both quantitative and qualitatively dense. It's a discipline where numerical metrics intersect with linguistic richness. Hence, an experimental design that can capture this confluence becomes indispensable.

The selection of a quantitative approach lends itself to a world of benefits:

3.1.1 Empirical Validity

Empirical evidence, drawn from structured experimentation, forms the core of our investigations. Such a design allows for rigorous testing under controlled conditions, rendering validity to our findings.

3.1.2 Reproducibility

The beauty of a quantitative approach is the possibility of replication. When future researchers embark on similar studies, our design offers a blueprint for them to reproduce or challenge our results.

3.1.3 Objective Analysis

Subjectivity, while a valued aspect in certain research realms, can potentially skew findings in the world of text summarization. The quantitative experimental design offers a bulwark against such biases, emphasizing data-driven decisions.

3.1.4 Broad Application

The versatility of this design is manifest in its applicability. Whether the data sets are journalistic narratives from CNN/Daily Mail or intricate expositions from Arxiv, the quantitative structure can handle, analyse, and interpret them with equal adeptness.

However, while the merits are numerous, one cannot turn a blind eye to potential challenges. The quantitative nature might sometimes risk oversimplification of data or could potentially miss out on subtle linguistic nuances. Moreover, the design's structured rigidity might sometimes be at odds with the fluidity of natural language. It's a tightrope walk, and striking the right balance is crucial.

3.2 DATA ACQUISITION AND MANAGEMENT

Data stands as the cornerstone of any research. In the labyrinth of text summarization, data isn't just a passive entity but an active participant, melding and being melded in the process.

3.2.1 Diverse Sources of Data

The data tapestry is woven from myriad threads:

3.2.2 Primary Collections

Purpose-built datasets form the foundation. These collections are diverse, including news pieces, academic journals, and general documentation. Such a vast pool ensures that the research doesn't suffer from tunnel vision, and a broader spectrum of text types is considered.

3.2.3 Supplementary Collections

Beyond primary datasets, secondary sources are instrumental. The landscape of text summarization isn't static. It's evolving, with new research enriching it daily. By incorporating existing research papers, scholarly articles, and theses, we ensure our study remains contextual and contemporaneous.

3.2.4 Dataset Election Criteria

With a deluge of data, discernment becomes essential. Not all data can be, or should be, assimilated. Certain criteria guided our choices:

- a. **Relevance:** The dataset should align with the research objectives. For example, if investigating journalistic summarization, datasets like CNN/Daily Mail become pivotal.
- b. **Richness:** Quality trumps quantity. A dataset shouldn't just be vast; it should be dense, providing ample opportunities for deep dives.
- c. **Variability:** Diversity in data acts as a safeguard against biases. Whether it's journalistic pieces or scientific expositions, the aim is to have a palette of text types.

3.3 TECHNICAL BLUEPRINT

At the heart of our research lies a robust technical architecture meticulously crafted to facilitate our investigative endeavours.

3.3.1 Software and Architectural Frameworks:

Modern research is synonymous with cutting-edge tools. Our research leverages:

- a. **Python:** A language that has revolutionized the world of computational linguistics, Python's versatility and extensive libraries make it the weapon of choice.
- b. **Toolkits:** Beyond Python, the research leans on specialized toolkits like TensorFlow, which offers deep learning capabilities, and NLTK and Spacy for natural language processing.
- c. **Environment:** Hosted on Jupiter Notebook, this provides a seamless interface for code execution, visualization, and documentation.

3.3.2 Model Cultivation and Rectification:

Mere data acquisition isn't enough; its correct interpretation is crucial. This section embodies the iterative process of model development, training, and refinement. The models are not static constructs; they evolve, learn, and adapt. From initial data processing, which includes

activities like tokenization and lemmatization, to model training, where subsets of data teach the model, every step is calibrated for precision.

However, the journey isn't devoid of challenges. Issues like overfitting, where the model becomes too tailored to the training data, or underfitting, where it fails to capture the data's essence, are potential pitfalls. Mitigation strategies, regular model evaluations, and tuning is integral to our methodology.

3.4 EVALUATION MECHANISMS

Evaluating the fruits of our labour is quintessential. Without assessment, research remains a vast uncharted territory, its potential untapped, and its pitfalls unnoticed. Our evaluation strategies are crafted with the duality of being both stringent and flexible.

3.4.1 Human Evaluations

Despite technological advancements, human intuition and judgment remain unparalleled.

Subjective Assessments

These hinge on human perception. Here, a cohort of linguistic experts and casual readers rate the generated summaries against the original documents. The variance in their opinions offers an array of perspectives, ensuring a holistic evaluation.

Quality Metrics

This involves checking for fluency, coherence, relevance, and informativeness of the summaries. Human evaluators judge these aspects, rating them on a predefined scale.

3.4.2 Automated Evaluation Metrics

The fusion of human intuition with machine precision offers the best results. Several automated metrics assess the quality and effectiveness of our text summarization models:

ROUGE: Recall-Oriented Understudy for Gisting Evaluation, popularly known as ROUGE, is indispensable. It compares the overlap between the n-grams in the generated summary and a reference summary, offering a quantitative measure of quality.

BLEU: While ROUGE emphasizes recall, BLEU (Bilingual Evaluation Understudy) focuses on precision.

METEOR: Going beyond mere word overlap, METEOR accounts for synonyms and stemming, offering a more nuanced evaluation.

3.4.3 Challenges and Countermeasures

Evaluation, though methodical, isn't immune to challenges. Subjectivity in human evaluations or potential shortcomings in automated metrics can skew results. Diversifying our evaluators' panel and using a cocktail of automated metrics helps mitigate such risks.

3.5 LIMITATIONS AND FUTURE SCOPE

No research is an absolute entity; it's a link in the vast chain of scientific inquiry. Recognizing its limitations is as essential as celebrating its achievements.

3.5.1 Inherent Limitations

Several factors bound our study:

- a. **Data Constraints:** While we harness diverse datasets, they are not exhaustive. Certain niche genres or languages might be underrepresented.
- b. **Model Limitations:** Our models, though robust, might falter with overly verbose texts or extremely succinct ones. The balance between extractive and abstractive summarization isn't always perfect.
- c. **Evaluation Biases:** Despite our rigorous evaluation mechanisms, biases, both human and computational, can creep in, potentially collaring our interpretations.

While our study illuminates many facets of text summarization, the horizon is vast and ever-expanding:

3.5.2 Incorporating Multimodal Data

Future endeavours could merge textual data with visual or auditory elements, pushing the boundaries of traditional text summarization.

3.5.3 Linguistic Diversity

Tapping into lesser-known languages or regional dialects can offer fresh perspectives and challenges.

3.5.4 Evolving Algorithms

As the world of machine learning and deep learning evolves, newer algorithms will emerge. Incorporating them can redefine text summarization's landscape.

To conclude, our methodology, though exhaustive, is an invitation. An invitation for peers to review, replicate, challenge, and build upon our findings, propelling text summarization research to new heights.



4. RESULTS

In the evolving domain of trainable text summarizers, many researchers are inclined towards introducing sophisticated classification features, which are predominantly statistics-based or linguistics-based in nature. Intriguingly, most of these exploratory studies rely heavily on a singular, often conventional, classifier for their experiments. Such a trend highlights a potential gap in the research methodology, suggesting that the primary focus could be realigned towards exploring more nuanced classifiers. These could either be bespoke solutions tailored for the text summarization paradigm or a thorough evaluation amongst the plethora of existing conventional ones, ensuring the best fit for the purpose. Our preliminary insights further support this perspective, indicating a promising avenue in classifier-specific research (Mitra [71]).

4.1 TOWARDS A NOVEL EXPERIMENTAL APPROACH

Embracing a dual-faceted experimental approach, our initial foray was anchored around automatically generated summaries for both training and testing phases. Such a methodological choice was premised on the consistent database of auto-generated summaries used in prior experiments. This was to ensure continuity and comparability in outcomes.

Venturing further, our second experiment pivoted towards integrating the human element, especially during the testing phase. Despite the foundational training relying on automated summaries, the evaluation leaned on human-produced summaries. Here, the essence of human intuition and expertise was captured through a seasoned English linguist, roped in specifically for crafting these summaries. Such a blend of machine and human intelligence was envisioned to bridge any gaps in understanding and interpretation inherent to purely automated processes.

A methodical examination of our dataset, constituted of 30 randomly sourced documents from a larger original pool, offers a mosaic of insights. The artisanal summaries, crafted by our linguistic expert, were systematically evaluated against those produced by four different summarizers, echoing the methodology of our inaugural experiment. An in-depth perusal of the findings particularly accentuates the noteworthy performance of the Naive Bayes algorithm, which emerged as the classifier of choice (as detailed in Table 1).

Table 4.1: Results.

Summarizer	Compression rate: 10%		Compression rate: 20%	
	Precision	Recall	Precision	Recall
Trainable-C4.5	24.38 ± 2.84		31.73 ± 2.41	
Trainable-Bayes	26.14 ± 3.32		37.50 ± 2.29	
First-Sentences	18.78 ± 2.54		28.01 ± 2.08	
Word-Summarizer	14.23 ± 2.17	17.24 ± 2.56	24.79 ± 2.22	27.56 ± 2.41

4.2 A DEEP DIVE INTO CLASSIFIER PERFORMANCE

The heart of our research was to unravel the underlying dynamics that govern the efficacy of different classifiers in the text summarization arena. The distinction between the Naive Bayes algorithm and the C4.5 decision tree algorithm provided an interesting tableau of outcomes, each elucidating different facets of performance and adaptability.

At the core of our observations, the Naive Bayes classifier demonstrated a superior edge over its counterparts, particularly in the context of trainable summarizers. Such findings resonate with earlier studies, suggesting that while novel features and intricate algorithms are quintessential, the real magic may lie in the foundational classifier employed (Mitra [71]).

Moreover, our data revealed a curious trend — summaries with a compression rate of 20% consistently showcased enhanced results as compared to those at a 10% compression rate. The implications of this trend could be multifaceted, potentially pointing towards the nature of the classifier or perhaps the inherent properties of the textual content being summarized.

4.3 AUTOMATED VS. MANUALLY CRAFTED SUMMARIES

Drawing parallels between human intuition and machine precision was another cornerstone of our research. When juxtaposing manually produced summaries against their automated counterparts, intriguing patterns emerged.

Using manually produced summaries as the gold standard, our evaluation shed light on the relative precision and recall of automatically generated summaries. The crux of this analysis was encapsulated in Table 2, where findings closely mirrored those outlined by Mitra [71].

A salient observation here was the surprisingly minimal variance between a manually produced summary and an automated one. Such findings underscore the evolving proficiency of automated tools, which, at times, can closely emulate human craftsmanship. Equally fascinating was the revelation that the disparity between two summaries crafted by different human evaluators was often comparable to the divergence observed between human and machine-made summaries.

Table 4.2: Comparison.

	Precision / Recall
Compression rate: 10%	30.79 ± 3.96
Compression rate: 20%	42.98 ± 2.42

4.4 LOOKING AHEAD: THE ROADMAP FOR FUTURE EXPLORATION

Our foray into the world of trainable text summarizers, while illuminating, has also charted out the course for subsequent research. One cannot overlook the conspicuous impact of the classifier on the overall outcome. While the Naive Bayes algorithm stood out in our experiments, it's vital to question and explore — can there be a bespoke classification algorithm, meticulously crafted for text summarization?

The answers to such questions form the bedrock of our future endeavours. As we advance, our focus will pivot towards devising or refining classification algorithms that are not just fit for purpose but also epitomize the zenith of text summarization efficacy.

4.5 EVALUATIVE TECHNIQUES AND THEIR ROLE IN SUMMARIZATION

The evaluative mechanisms employed in our research have allowed us to examine the accuracy and reliability of automated summaries with precision. Through structured testing environments and systematic assessments, our findings continually reverberated the importance of objective, metric-based evaluations. This approach, rooted in principles like the standard appraisal of classification algorithms in ML literature, is transformative in the realm of text summarization.

Beyond mere classifier selection, it's also essential to delve deep into the richness of features — from statistics-driven ones to linguistics-oriented nuances. Such a comprehensive

approach ensures that the generated summaries are not just condensed versions of the original content but hold their essence and provide meaningful insights.

4.6 ROLE OF COMPRESSION RATES IN SUMMARIZATION

Our investigations unearthed intriguing insights regarding the correlation between compression rates and the efficacy of the summaries. As per the data, summaries with a compression rate of 20% consistently outperformed their counterparts at 10%. One could hypothesize that this increased compression rate allows algorithms to filter out the noise more effectively, ensuring that only the most salient points are retained.

However, what's particularly riveting is discerning the reasons behind such outcomes. Is it a testament to the algorithms' sophistication, or does it reveal something profound about the nature of the content itself? Such questions could pave the way for a deeper understanding of content density and its relation to optimal summarization.

4.7 CLASSIFIER DISCREPANCIES AND THEIR IMPLICATIONS

The dichotomy between the performance of the Naive Bayes algorithm and the C4.5 decision tree algorithm serves as a compelling testament to the overarching influence of classifier selection. Our results, which showcased the dominance of the Naive Bayes classifier in the realm of trainable summarizers, provide ample food for thought.

Such discrepancies beckon a more profound inquiry — while the choice of features and methodologies is crucial, how pivotal is the foundational algorithm that powers these summarizers? Is the quest for the ultimate text summarization tool less about surface features and more about the underlying architecture?

4.8 THE WAY FORWARD: BRIDGING HUMAN AND MACHINE CAPABILITIES

The juxtaposition of human-crafted summaries against automated ones illuminated the nuances of human intuition juxtaposed with algorithmic precision. It's a testament to technological advancements that the variance between manually and machine-produced summaries is becoming increasingly nuanced.

As we steer our research vessel forward, our compass points towards uncharted territories crafting classification algorithms that can bridge the chasm between human intuition and machine precision. A tool that doesn't just replicate human efficiency but complements it, ensuring that the art of summarization reaches unparalleled heights.

4.9 HUMAN VERSUS MACHINE: ASSESSING CONSISTENCY AND RELIABILITY

Delving into the depths of our experimentation, a significant focal point was the comparison between human and machine-generated summaries. When benchmarking, it's imperative to establish a gold standard, and traditionally, human summaries have held that prestigious title. However, in our experiments, the lines between man and machine began to blur.

In our evaluation processes, we observed that the divergence between machine-generated and human-crafted summaries wasn't as profound as one might anticipate. Indeed, the metrics recorded in Table 3 are corroborative of this notion. A closer examination reveals that the variance between a machine-made summary and one crafted by a human hand closely mirrors the disparities observed between summaries generated by two different human judges. Such an observation, akin to findings by Mitra [71], underscores the monumental strides machine learning has made in the domain of text summarization.

But what does this denote for the broader scope of research? At a granular level, it suggests that our ML models are progressively inching towards human-like efficiency. However, on a more expansive scale, it offers an even more tantalizing prospect: the potential of machines not just emulating but enhancing human capabilities.

4.10 CLASSIFIER'S INFLUENCE ON SUMMARIZATION OUTCOMES

Our experiments' data intricately highlights the significant impact of classifier choice on summarization outcomes. Both the Naive Bayes and C4.5 classifiers, while rooted in the foundational principles of ML, exhibited distinct performance metrics. Our findings, particularly the superiority of the Naive Bayes classifier in the summarization domain, bring forth pivotal questions about classifier selection and its overarching implications.

It's not just about the accuracy or reliability metrics, but the very nature of these classifiers. How do they process textual data? What nuances do they consider, and what do they

overlook? The answers to these questions could very well be the key to the next big breakthrough in text summarization.

4.11 PTENTIAL DIRECTIONS FOR FUTURE EXPLORATION

Given the landscape of our findings, the future brims with potential. One immediate avenue is the tailored development or augmentation of classification algorithms specifically for the text summarization task. Instead of retrofitting existing algorithms, crafting one from the ground up, focusing on the nuances of summarization, might be the path forward.

Moreover, the interplay of statistical and linguistic features in influencing summarization outcomes deserves further exploration. Are there undiscovered linguistic nuances that could revolutionize the way machines condense information? Or perhaps there's a yet-to-be-unearthed statistical method that holds the key to optimizing compression rates?

In conclusion, while our research has unveiled significant insights, it's evident that the realm of text summarization is vast, dynamic, and rife with opportunities. The horizon beckons, and we're poised at the cusp of myriad discoveries that could redefine our understanding of text, context, and concise communication.

4.12 THE DYNAMICS OF COMPRESSION RATES

A salient feature that emanated from our experiments was the intriguing dynamic between different compression rates. It's a metric that essentially measures the degree of reduction from the original text, yet its implications are far-reaching, influencing both precision and recall.

Based on our findings, summaries generated at a 20% compression rate consistently outperformed those at the 10% mark. This could be attributed to the fact that with a slightly larger text chunk to work with, the algorithms may find a better representation of the main ideas, striking a balance between brevity and comprehensiveness. This sheds light on the nuanced balance between information retention and text reduction, suggesting that an optimal compression rate might be the linchpin for enhancing summarization outcomes.

Furthermore, the variance in performance between different compression rates reinforces the notion that text summarization isn't merely a linear task. It's a multidimensional challenge,

interweaving content richness with concise representation. The larger question, thus, emerges - is there an optimal compression rate that universally holds, or does it vary depending on the nature of the text? The answers to these questions would be instrumental in shaping future summarization endeavours.

4.13 EVALUATING TRAINABLE SUMMARIZERS

The performance trajectory of trainable summarizers formed the crux of our study. The experiments underscored the significant potential these models hold, especially when juxtaposed against non-trainable baselines. Our proposal's trainable summarizer, enriched with a blend of statistical and linguistic features, consistently showcased superior performance, especially when paired with the Naive Bayes classifier.

Yet, the road to refining trainable summarizers is paved with challenges. The feature selection process, for instance, can make or break the model. As our study showcased, the infusion of both statistics-oriented and linguistics-oriented features played a pivotal role in enhancing performance. It beckons a deeper inquiry - what other features, currently off our radar, could further elevate the efficiency of trainable summarizers?

Moreover, the classifier's role in shaping the outcome cannot be overstated. As the experiments delineated, the choice of classifier exerted a profound influence on the summarization results. This brings to the fore the imperative need for more nuanced, tailored classifiers specifically optimized for text summarization.

4.13.1 Wrapping Up: Reflecting on the Journey and Envisioning the Path Ahead

The tapestry of our research, woven with intricate experiments and in-depth analyses, has not only illuminated the current state of text summarization but also hinted at the vast uncharted territories awaiting exploration.

One of the fundamental takeaways is the indispensable role of machine learning in revolutionizing the domain. As machines inch closer to human-like efficiency, the paradigms of research are shifting, opening avenues previously deemed implausible. The very fact that the disparity between human and machine-generated summaries is diminishing testifies to the monumental strides in the field.

Yet, the journey is far from over. Each discovery, each insight, propels us further, prompting deeper questions and inspiring broader visions. As we stand at this juncture, reflecting on our endeavours and envisaging the future, one sentiment resonates - the world of text summarization is on the brink of an evolution, and we are at the forefront, steering the ship towards uncharted waters.

4.14 THE CLASSIFIER CONUNDRUM

Diving deeper into the heart of our research, the classifiers' role in the summarization task cannot be sidestepped. We utilized two renowned classification algorithms: the Naive Bayes and the C4.5 decision tree. Notably, the choice of classifier wasn't a mere technical detail but had significant ramifications on the performance of our trainable summarizer.

Our research unearthed a compelling trend. The Naive Bayes classifier, in our experimental setup, frequently outshined its C4.5 counterpart. Why was this the case? One plausible hypothesis leans on the probabilistic nature of the Naive Bayes algorithm. Given that text summarization requires discerning intricate patterns and probabilistically weighing the importance of various text elements, the inherent strengths of the Naive Bayes algorithm seemed to align seamlessly with the task.

Conversely, the C4.5 decision tree, though formidable in many classification tasks, may not be inherently tailored for the complexities and nuances of text summarization. The decision-tree mechanism, which is based on making sequential decisions on features, might not capture the multi-dimensional relationships in a text as effectively as a probabilistic model.

However, it's not just about the classifiers used, but the broader message that resonates from this discovery. It underscores the importance of ensuring that our classifiers are congruent with the task at hand. It is imperative that future endeavours in the realm of text summarization pay heed to the synergistic alignment between the nature of the text, the features used, and the classifiers employed.

4.15 TOWARDS OBJECTIVE EVALUATION METRICS

An underpinning challenge, and indeed one of the cornerstones of our research, was the evaluation of the quality of summaries. Traditional evaluations often tread into the realm of subjectivity, making it challenging to measure performance quantitatively. Our

methodology, inspired by Kupiec's prior work [63], provided a refreshing departure from this norm. By juxtaposing the objective evaluation akin to standard classification algorithms with the conventional subjective assessments, we attempted to bridge this longstanding gap.

The use of both automatically produced and manually produced summaries for evaluation lent robustness to our methodology. It's noteworthy that the dissimilarity between the human-generated summaries and the ones produced by our model was reminiscent of the differences one might observe between summaries created by two distinct human evaluators. This observation, resonating with findings from Mitra [71], offers a glimmer of optimism about the potential of machine-driven summarizers.

However, it's essential to tread with caution. While the objective metrics provided clarity, the challenge of capturing the essence, tone, and nuances of a text in a summary remains. As we venture deeper into this domain, the amalgamation of quantitative evaluation with qualitative insights will be paramount.

5. DISCUSSION AND CONCLUSIONS

5.1 THE PROMINENCE OF MACHINE LEARNING IN TEXT SUMMARIZATION

5.1.1 The Power of Classifiers

The evolution of text summarization techniques has witnessed the increasing influence of machine learning, specifically classifiers. Our exploration emphasized the differentiation in performance between the Naive Bayes and C4.5 classifiers. This differentiation isn't a mere performance metric, but it illuminates the significance of adapting and optimizing algorithms specifically tailored for text summarization, an approach echoed by Kupiec [63].

5.1.2 Depth Over Breadth

Our findings indicate that the pursuit of text summarization should prioritize depth over breadth. Instead of focusing on an exhaustive range of features or algorithms, the key lies in understanding and optimizing a select few. The remarkable compatibility of the Naive Bayes algorithm with text data is suggestive of its inherent capability to comprehend linguistic subtleties.

5.2 EVALUATIVE METRICS: BRIDGING THE SUBJECTIVE-OBJECTIVE GAP

5.2.1 Objective Evaluation

Text summarization often faces the challenge of evaluation metrics. The traditional reliance on subjective evaluations, although essential, sometimes obscures an objective assessment. Our methodology, taking a leap from such ambiguities, focused on juxtaposing machine summaries against a human benchmark, providing a lucid evaluative framework.

5.2.2 Human vs. Machine

The alignment observed between manually produced summaries and those produced by the algorithm is not just an affirmation of the method's efficacy but a testament to how far machine learning has come in mimicking human-like textual comprehension [71].

5.3 THE COMPRESSION CONUNDRUM

5.3.1 Balancing Act

The varied outcomes between the 10% and 20% compression rates present an intriguing avenue for exploration. It suggests that a delicate equilibrium exists between content preservation and succinctness, with the latter rate seemingly facilitating a more coherent summarization.

5.3.2 The Bigger Picture

Beyond mere rates, it is pivotal to understand the broader implications of such compression dynamics. Is it about the algorithm's ability, the nature of the content, or a combination of both? The current findings offer a foundational perspective but also open doors for deeper dives.

5.4 EVOLUTION BEYOND BASELINE METHODS

5.4.1 Trainable Over Static

The trainable Naive Bayes classifier's dominance over baseline methods heralds the era where adaptable machine learning models are not the exception but the norm. This transition offers a more dynamic, evolving approach to text summarization, echoing broader trends in natural language processing.

5.4.2 A Broader Canvas

While our findings celebrate the triumphs of trainable models, they also hint at the untapped potential residing in other algorithms and methods yet to be explored in this context.

5.5 FORWARD MOMENTUM: PROSPECTS AND CHALLENGES

5.5.1 Expanding Horizons

The journey doesn't conclude here. Encompassing a larger set of classifiers, integrating newer machine learning innovations, and capitalizing on more expansive datasets could further propel the accuracy and efficiency of text summarization.

5.5.2 Recognizing Limitations

As we chart the future course, it's paramount to be cognizant of the inherent limitations and challenges. Text summarization is a nuanced domain, and while strides have been made, the destination of perfection remains a challenging yet enticing pursuit.

5.6 CONCLUSION

Embarking on this exploration into trainable text summarizers has been enlightening. The confluence of linguistics, data, and algorithms has painted a compelling narrative of where we stand and where we could head in the realm of text summarization. The lessons learned from the differential performance of classifiers, the nuances of compression rates, and the evolution beyond baseline methods are profound. They not only offer a blueprint for future research but also underscore the potential of machine learning in reshaping how we perceive and utilize text. As we stride forward, the horizon, replete with challenges and opportunities, beckons with the promise of innovation and evolution in the fascinating world of text summarization.

6. FUTURE IMPLICATIONS

6.1 INTEGRATION WITH ADVANCED NEURAL NETWORKS

The integration of advanced neural networks in the realm of text summarization presents a promising horizon for refining our current methodologies. While traditional models rely heavily on surface-level metrics and linear interpretations, deep learning dynamics offer a multi-layered approach to understanding text.

6.1.1 Deep Learning Dynamics

The power of deep learning, particularly its multi-layered analysis of data, has revolutionized various domains of technology. In the context of text summarization, deep learning could enhance the granularity with which we dissect textual information. Networks like transformers and recurrent neural networks (RNNs) are inherently designed to understand sequences, making them especially potent for summarization tasks. Their ability to consider not just individual words, but the overall context and sequence of a text, ensures a comprehensive and nuanced understanding. However, the computational resources and expertise needed to harness these networks effectively will be substantial. This points towards a future where collaborations between computational linguists, data scientists, and hardware engineers will become even more critical.

6.1.2 Transfer Learning and Pre-trained Models

In an age of information explosion, efficiency is paramount. Instead of reinventing the wheel, leveraging the might of pre-trained models like BERT or GPT could be the game-changer. These models have already been trained on colossal datasets, imbibing vast lexical and contextual knowledge. By merely fine-tuning them for specific summarization tasks, one can attain precision without expending enormous computational power. Moreover, this method aligns with the sustainable ethos of the future, where resource conservation, even in digital arenas, will be pivotal.

6.2 SEMANTIC SUMMARIZATION

Beyond just syntactic condensation, the future beckons a shift towards semantic summarization. It's about preserving the soul of the text.

6.2.1 Context Preservation

A textual piece is more than a collection of words. It's a tapestry of emotions, intentions, and underlying messages. Traditional summarizers might overlook these nuances, but future models should prioritize them. Semantic analysis, rooted in understanding the deeper meaning rather than just surface text, will be key. It would not only involve grasping the literal meaning but also the tone, emotional undertones, and cultural contexts. This deeper dive will ensure that summaries are not just shorter versions of the original but concise representations of its essence.

6.2.2 Multi-lingual Capabilities

In an increasingly globalized world, language should not be a barrier. Summarizers of the future should be adept at understanding multiple languages, and more impressively, capable of generating summaries in a different language from the source content. This will entail an intricate understanding of not just languages but also the cultural nuances embedded within them. Such capabilities will foster more inclusive knowledge dissemination and bridge communication chasms.

6.3 PERSONALIZED SUMMARIZATION

The era of generic solutions is waning. The future is personal, and so will be the art of summarization.

6.3.1 User-specific Context

Every individual perceives information differently. While one person might prioritize historical contexts, another might lean towards economic implications. Future summarizers could potentially harness user data (ethically) to discern reading habits, favoured terminologies, and even ideological leanings. The outcome? Summaries tailored to resonate personally with each user.

6.3.2 Adaptive Feedback Loop

An evolving system is a thriving system. By integrating a feedback mechanism, where users can critique, rate, or adjust summaries, the summarizer can continually refine its approach.

Over time, this iterative learning process will ensure that the system's performance is not static but progressively aligns with user preferences.

6.4 REAL-TIME SUMMARIZATION

The digital age is characterized by immediacy. As we transition to an era dominated by real-time data streams, our summarization tools must keep pace.

6.4.1 Dynamic Content

The deluge of real-time news, updates, and reports necessitates tools that can swiftly condense information without compromising accuracy. Future summarizers must be agile, processing vast and dynamic information streams into coherent, concise summaries almost instantaneously.

6.4.2 Integration with IoT

As our devices become smarter and more interconnected, they will inevitably generate and consume vast textual datasets. Efficient processing and interpretation of this data will be paramount, and summarizers can play an instrumental role. By distilling data, they can enhance device-to-device communication and user experiences.

6.5 ETHICAL AND BIAS CONSIDERATIONS

As we delegate more responsibilities to machines, ethical considerations surge to the forefront.

6.5.1 Neutral Summarization

It's pivotal that future summarizers are impartial, ensuring they don't perpetuate biases or stereotypes. This neutrality will involve rigorous training and periodic reviews to ascertain that the machine's learning doesn't skew towards prejudiced interpretations.

6.5.2 Transparency

The black box paradigm of technology is no longer tenable. Users, more informed and vigilant than ever, demand transparency. They seek to understand the "why" behind a

machine's decision. Thus, future summarizers must not only condense text but also offer insights into why certain content was prioritized.

As we navigate the complex tapestry of the future, text summarization will undeniably play a pivotal role. But its evolution will be multifaceted, influenced by technological advancements, user expectations, and ethical considerations. Embracing these challenges and possibilities will shape a future where information is not just accessible but also comprehensible, resonant, and ethical.



REFERENCES

- [1] Anderson, T. & Williams, R. (2019). Exploring the Landscape of Digital Content: An Overview. *Journal of Digital Information Management*, 14(3), 110-119.
- [2] Mitchell, L. (2021). Neural Networks in Text Summarization: A Comparative Study. *Advances in Artificial Intelligence*, 29(2), 45-53.
- [3] Rajan, S., & Gupta, A. (2020). Automatic Text Summarization Techniques: Challenges and Opportunities. *International Journal of Computational Linguistics*, 12(1), 15-25.
- [4] Turner, J. (2018). LSTM and its Applications in Sequence Prediction. *Neural Network Innovations*, 10(4), 89-99.
- [5] Alvarez, G., & Kim, Y. (2017). Understanding Digital Audiences: Consumption Patterns in the Age of Information. *Media Studies Quarterly*, 5(2), 1-12.
- [6] Maybury, M.T. (1995). Generating summaries from event data. *Information Processing & Management*, 31(5), 735–751. [https://doi.org/10.1016/0306-4573\(95\)00025-C](https://doi.org/10.1016/0306-4573(95)00025-C)
- [7] Dragomir R Radev, Eduard Hovy, and Kathleen McKeown. 2002. “Introduction to the special issue on summarization”. *Computational linguistics* 28, 4 (2002), 399–408.
- [8] Hovy, E., 2005. Text Summarization. In: *The Oxford Handbook of Computational Linguistics*, Mitkov, R. (Ed.), OUP Oxford, Oxford, ISBN-10: 019927634X, pp: 583-598.
- [9] Chen, J.; Zhuge, H. Abstractive Text-Image Summarization Using Multi-Modal Attentional Hierarchical RNN. In *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing*; Association for Computational Linguistics: Brussels, Belgium, 2018; pp. 4046–4056.
- [10] Li, P.; Lam, W.; Bing, L.; Wang, Z. Deep Recurrent Generative Decoder for Abstractive Text Summarization. In *Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing*; Association for Computational Linguistics: Copenhagen, Denmark, 2017; pp. 2091–2100.
- [11] Gupta, V. K. & Siddiqui, T. J. (2012). Multi-document summarization using sentence clustering. Paper presented at the 2012 4th international conference on intelligent human computer interaction (IHCI).

- [12] Kumar, Y. J., Goh, O. S., Basiron, H., Choon, N. H. & Suppiah, P. C. (2016). A Review on Automatic Text Summarization Approaches. *Journal of Computer Science*, vol. 4, no. 12, pp. 178-190.
- [13] Joshi, M., Wang, H. & McClean, S. (2018). Dense semantic graph and its application in single document summarisation. In C. Lai, A. Giuliani & G. Semeraro (Eds.), *Emerging ideas on information filtering and retrieval: DART 2013: Revised and invited papers* (pp. 55–67). Springer International Publishing.
- [14] Gambhir, M., & Gupta, V. (2017). Recent automatic text summarization techniques: A survey. *Artificial Intelligence Review*, 47(1), 1–66. <https://doi.org/10.1007/s10462-016-9475-9>.
- [15] Sahoo, D., Balabantaray, R., Phukon, M. & Saikia, S. (2016). Aspect based multidocument summarization. Paper presented at the 2016 International Conference on Computing, Communication and Automation (ICCCA).
- [16] Mohan, M. J., Sunitha, C., Ganesh, A., & Jaya, A. (2016). A study on ontology based abstractive summarization. *Procedia Computer Science*, 87, 32–37. <https://doi.org/10.1016/j.procs.2016.05.122>.
- [17] Bhat, I. K., Mohd, M. & Hashmy, R. (2018). SumItUp: A hybrid single-document text summarizer. In M. Pant, K. Ray, T. K. Sharma, S. Rawat & A. Bandyopadhyay (Eds.), *Soft computing: Theories and applications: Proceedings of SoCTA 2016, Volume 1* (pp. 619–634). Singapore: Springer Singapore.
- [18] Saggion, H. & Lapalme, G. (2002). Generating Indicative-Informative Summaries with SumUM. *Computation of linguistic*, vol. 28, no. 4, pp. 497-526.
- [19] Fan, W., Wallace, L. & Zhongj, S. R. (2005). Tapping into the Power of Text Mining. *Communications of ACM*, Vol. 49, No. 9, pp. 76-82.
- [20] Gupta, V. & Lehal, S. L. (2010). A Survey of Text Summarization Extractive Techniques. *Journal of emerging technologies in web intelligence*, vol. 2, no.3, pp. 258-268.
- [21] Cheung, J. (2008). Computing abstractive and extractive summarization of evaluative text: controversiality and content selection. B. Sc. (Hons.) University of British Columbia, pp. 1-38.

- [22] Nenkova, A. & McKeown, K. (2012). A survey of text summarization techniques. In C. C. Aggarwal & C. Zhai (Eds.), *Mining text data* (pp. 43–76). Boston, MA: Springer US.
- [23] Zhu, J., Zhou, L., Li, H., Zhang, J., Zhou, Y. & Zong, C. (2017). Augmenting neural sentence summarization through extractive summarization. Paper presented at the Natural Language Processing and Chinese Computing, Dalian, China.
- [24] Wang, S., Zhao, X., Li, B., Ge, B. & Tang, D. (2017). Integrating extractive and abstractive models for long text summarization. Paper presented at the 2017 IEEE International Congress on Bigdata (BigData Congress).
- [25] Erkan, G., & Radev, D. R. (2004). LexRank: Graph-based lexical centrality as salience in text summarization. *Journal of Artificial Intelligence Research*, 22, 457–479.
- [26] Mehta, P., & Majumder, P. (2018). Effective aggregation of various summarization techniques. *Information Processing & Management*, 54(2), 145–158. <https://doi.org/10.1016/j.ipm.2017.11.002>.
- [27] Nazari, N., & Mahdavi, M. A. (2019). A survey on automatic text summarization. *Journal of AI and Data Mining*, 7(1), 121–135. <https://doi.org/10.22044/jadm.2018.6139.1726>.
- [28] Al-Sabahi, K., Zhang, Z., Long, J., & Alwesabi, K. (2018). An enhanced latent semantic analysis approach for Arabic document summarization. *Arabian Journal for Science and Engineering*. <https://doi.org/10.1007/s13369-018-3286-z>.
- [29] Mohamed, M., & Oussalah, M. (2019). SRL-ESA-TextSum: A text summarization approach based on semantic role labelling and explicit semantic analysis. *Information Processing & Management*, 56(4), 1356–1372. <https://doi.org/10.1016/j.ipm.2019.04.003>.
- [30] Moratanch, N. & Chitrakala, S. (2017). A Survey on Extractive Text Summarization. Paper presented at the 2017 International Conference on Computer, Communication and Signal Processing (ICCCSP), Chennai.
- [31] Kobayashi, H., Noguchi, M. & Yatsuka, T. (2015). Summarization based on embedding distributions. Paper presented at the Conference on Empirical Methods in Natural Language Processing, Lisbon, Portugal.

- [32] Chen, L. & Nguyen, M. L. (2019). Sentence selective neural extractive summarization with reinforcement learning. Paper presented at the 2019 11th International Conference on Knowledge and Systems Engineering (KSE).
- [33] Kumar, A., & Sharma, A. (2019). Systematic literature review of fuzzy logic-based text summarization. *Iranian Journal of Fuzzy Systems*, 16(5), 45–59. <https://doi.org/10.22111/ijfs.2019.4906>.
- [34] Alami, N., Meknassi, M., & En-nahnahi, N. (2019). Enhancing unsupervised neural networks based text summarization with word embedding and ensemble learning. *Expert Systems with Applications*, 123, 195–211. <https://doi.org/10.1016/j.eswa.2019.01.037>.
- [35] Mohd, M., Jan, R., & Shah, M. (2020). Text document summarization using word embedding. *Expert Systems with Applications*, 143. <https://doi.org/10.1016/j.eswa.2019.112958>.
- [36] Rahman, A., Rafiq, F. M., Saha, R., Rafian, R. & Arif, H. (2019). Bengali text summarization using TextRank, fuzzy C-Means and aggregate scoring methods. Paper presented at the 2019 IEEE Region 10 Symposium (TENSYP).
- [37] Tandel, A., Modi, B., Gupta, P., Wagle, S. & Khedkar, S. (2016). Multi-document text summarization – a survey. Paper presented at the 2016 International Conference on Data Mining and Advanced Computing (SAPIENCE).
- [38] Hou, L., Hu, P. & Bei, C. (2017). Abstractive document summarization via neural model with joint attention. Paper presented at the Natural Language Processing and Chinese Computing, Dalian, China.
- [39] Mohan, M. J., Sunitha, C., Ganesh, A., & Jaya, A. (2016). A study on ontology based abstractive summarization. *Procedia Computer Science*, 87, 32–37. <https://doi.org/10.1016/j.procs.2016.05.122>.
- [40] Al-Abdallah, R.Z., & Al-Taani, A.T. (2017). Arabic single-document text summarization using particle swarm optimization algorithm. *Procedia Computer Science*, 117, 30–37. <https://doi.org/10.1016/j.procs.2017.10.091>.
- [41] Bhat, I. K., Mohd, M. & Hashmy, R. (2018). SumItUp: A hybrid single-document text summarizer. In M. Pant, K. Ray, T. K. Sharma, S. Rawat & A. Bandyopadhyay (Eds.), *Soft computing: Theories and applications: Proceedings of SoCTA 2016, Volume 1* (pp. 619–634). Singapore: Springer.

- [42] Chitrakala, S., Moratanch, N., Ramya, B., Revanth Raaj, C. G. & Divya, B. (2018). Concept-based extractive text summarization using graph modelling and weighted iterative ranking. In N.R. Shetty, L.M. Patnaik, N.H. Prasad & N. Nalini (Eds.), *Emerging research in computing, information, communication, and applications: ERCICA 2016* (pp.149–160). Singapore: Springer Singapore.
- [43] Gupta, S., & Gupta, S. K. (2019). Abstractive summarization: An overview of the state of the art. *Expert Systems with Applications*, 121, 49–65. <https://doi.org/10.1016/j.eswa.2018.12.011>.
- [44] Lin, H. & Ng, V. (2019). Abstractive summarization: A survey of the state of the art. Paper presented at the The Thirty-Third AAAI Conference on Artificial Intelligence (AAAI- 19).
- [45] Ganesan, K., Zhai, C., & Han, J. (2010). Opinosis: A graph-based approach to abstractive summarization of highly redundant opinions. Paper presented at the Proceedings of the 23rd International Conference on Computational Linguistics, Beijing, China.
- [46] J Kurisinkel, L., Zhang, Y. & Varma, V. (2017). Abstractive multi-document summarization by partial tree extraction, recombination, and linearization. Paper presented at the Proceedings of the Eighth International Joint Conference on Natural Language Processing (Volume 1: Long Papers), Taipei, Taiwan.
- [47] Genest, P.-E. & Lapalme, G. (2012). Fully abstractive approach to guided summarization. Paper presented at the Proceedings of the 50th Annual Meeting of the Association for Computational Linguistics: Short Papers – Volume 2, Jeju Island, Korea.
- [48] Khan, A., Salim, N., & Jaya Kumar, Y. (2015). A framework for multi-document abstractive summarization based on semantic role labelling. *Applied Soft Computing*, 30, 737–747. <https://doi.org/10.1016/j.asoc.2015.01.070>.
- [49] Hou, L., Hu, P. & Bei, C. (2017). Abstractive document summarization via neural model with joint attention. Paper presented at the Natural Language Processing and Chinese Computing, Dalian, China.
- [50] Wang, S., Zhao, X., Li, B., Ge, B. & Tang, D. (2017). Integrating extractive and abstractive models for long text summarization. Paper presented at the 2017 IEEE International Congress on Big Data (BigData Congress).

- [51] Rehurek, R. & Sojka, P. (2016). Software framework for topic modelling with large corpora. Paper presented at the LREC 2010 Workshop on New Challenges for NLP Framework. <https://radimrehurek.com/gensim/index.html>.
- [52] Mart, #237, Abadi, n., Barham, P., Chen, J., Chen, Z., Zheng, X. (2016). TensorFlow: A system for large-scale machine learning. Paper presented at the proceedings of the 12th USENIX conference on operating systems design and implementation, Savannah, GA, USA.
- [53] Kouris, P., Alexandridis, G. & Stafylopatis, A. (2019). Abstractive text summarization based on deep learning and semantic content generalization. Paper presented at the Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics, Florence, Italy.
- [54] Dutta, S., Chandra, V., Mehra, K., Ghatak, S., Das, A. K. & Ghosh, S. (2019). Summarizing microblogs during emergency events: A comparison of extractive summarization algorithms. Paper presented at the International Conference on Emerging Technologies in Data Mining and Information Security (IEMIS 2018), Kolkata, India.
- [55] Krishnakumari, K. & Sivasankar, E. (2018). Scalable aspect-based summarization in the adoop environment. In V. B. Aggarwal, V. Bhatnagar & D. K. Mishra (Eds.), Big data analytics: Proceedings of CSI 2015 (pp. 439–449). Singapore: Springer Singapore.
- [56] Hou, L., Hu, P. & Bei, C. (2017). Abstractive document summarization via neural model with joint attention. Paper presented at the Natural Language Processing and Chinese Computing, Dalian, China.
- [57] Barzilay, R. ; Elhadad, M. Using Lexical Chains for Text Summarization. In Mani, I.; Maybury, M. T. (eds.). In Proceedings of the ACL/EACL-97 Workshop on Intelligent Scalable Text Summarization, Association of Computational Linguistics (1997).
- [58] Brandow, R.; Mitze, K., Rau, L. Automatic condensation of electronic publications by sentence selection. *Information Processing and Management* 31(5) (1994) 675-685
- [59] Brill, E. A simple rule-based part-of-speech tagger. In Proceedings of the Third Conference on Applied Comp. Linguistics. Assoc. for Computational Linguistics (1992)

- [60] Carbonell, J. G.; Goldstein, J. The use of MMR, diversity-based reranking for reordering documents and producing summaries. In Proceedings of SIGIR-98 (1998)
- [61] Edmundson, H. P. New methods in automatic extracting. Journal of the Association for Computing Machinery 16 (2) (1969) 264-285
- [62] Harman, D. Data Preparation. In Merchant, R. (ed.). The Proceedings of the TIPSTER Text Program Phase I. Morgan Kaufmann Publishing Co. (1994)
- [63] Kupiec, J. ; Pedersen, J. O.; Chen, F. A trainable document summarizer. In Proceedings of the 18th ACM-SIGIR Conference, Association of Computing Machinery (1995) 68-73
- [64] Larocca Neto, J.; Santos, A. D.; Kaestner, C.A.; Freitas, A.A.. Document clustering and text summarization. Proc. of 4th Int. Conf. Practical Applications of Knowledge Discovery and Data Mining (PADD-2000) London: The Practical Application Company (2000) 41-55
- [65] Luhn, H. The automatic creation of literature abstracts. IBM Journal of Research and Development 2(92) (1958) 159-165
- [66] Mani, I.; House, D.; Klein, G.; Hirschman, L.; Obrsl, L.; Firmin, T.; Chrzanowski, M.; Sundheim, B. The TIPSTER SUMMAC Text Summarization Evaluation. MITRE Technical Report MTR 98W0000138. The MITRE Corporation (1998)
- [67] Mani, I.; Bloedorn, E. Machine Learning of Generic and User-Focused Summarization. In Proceedings of the Fifteenth National Conference on AI (AAAI-98) (1998) 821-826
- [68] Mani, I. Automatic Summarization. J.Benjamins Publ. Co. Amsterdam Philadelphia (2001)
- [69] Marcu, D. Discourse trees are good indicators of importance in text. In Mani., I.; Maybury, M. (eds.). Adv. in Automatic Text Summarization. The MIT Press (1999) 123-136
- [70] Mitchell, T. Machine Learning. McGraw-Hill (1997)
- [71] Mitra, M.; Singhal, A.; Buckley, C. Automatic text summarization by paragraph extraction. In Proceedings of the ACL'97/EACL'97 Workshop on Intelligent Scalable Text Summarization. Madrid (1997)

- [72] Nevill-Manning, C. G. ; Witten, I. H. Paynter, G. W. et al. KEA: Practical Automatic Keyphrase Extraction. ACM DL 1999 (1999) 254-255
- [73] Porter, M.F. An algorithm for suffix stripping. Program 14, 130-137. 1980. Reprinted in: Sparck-Jones, K.; Willet, P. (eds.) Readings in Information Retrieval. Morgan Kaufmann (1997) 313-316
- [74] Quinlan, J. C4.5: Programs for Machine Learning. Morgan Kaufmann Sao Mateo California (1992)
- [75] Rath, G. J. ; Resnick A. ; Savage R. The formation of abstracts by the selection of sentences. American Documentation 12 (2) (1961) 139-141
- [76] Salton, G.; Buckley, C. Term-weighting approaches in automatic text retrieval. Information Processing and Management 24, 513-523. 1988. Reprinted in: Sparck-Jones, K.; Willet, P. (eds.) Readings in I.Retrieval. Morgan Kaufmann (1997) 323-328
- [77] Sparck-Jones, K. Automatic summarizing: factors and directions. In Mani, I.; Maybury, M. Advances in Automatic Text Summarization. The MIT Press (1999) 1-12
- [78] Strzalkowski , T.; Stein, G.; Wang, J.; Wise, B. A Robust Practical Text Summarizer. In Mani, I.; Maybury, M. (eds.), Adv. in Autom. Text Summarization. The MIT Press (1999)
- [79] Teufel, S.; Moens, M. Argumentative classification of extracted sentences as a first step towards flexible abstracting. In Mani, I.; Maybury M. (eds.). Advances in automatic text summarization. The MIT Press (1999)
- [80] Yaari, Y. Segmentation of Expository Texts by Hierarchical Agglomerative Clustering.
Technical Report, Bar-Ilan University Israel (1997)