

KEY PHRASE-BASED AUTOMATIC TEXT DOCUMENT SUMMARIZATION

¹Tahoor Zaidi, ²Suraj Malik, ³Mukesh Rawat

^{1,2}Dept. of CSE, IIMT University, Meerut

³ Dept. of CSE, MIET, Meerut

¹ tahoorzaidi@gmail.com, ² surajmalik@iimtindia.net, ³
mukesh.rawat@miet.ac.in

Abstract

In NLP and information retrieval, summarization of text is an important function in summarizing long documents into readable, concise summaries. The process is especially applicable in the analysis of long texts like research papers, legal briefs, and news articles, where important information needs to be pulled out. The nature of summarization is to preserve meaning in the original text while reducing its length. This research focuses on two primary summarization methods: extractive and abstractive. Extractive methods pull out important sentences from the original text, while abstractive methods generate new sentences to summarize. The research focuses on three top algorithms: TextRank, Seq2Seq, and BART. TextRank, a graph-based algorithm, is best suited for short texts. Seq2Seq, a deep learning-based algorithm, produces highly accurate summaries by understanding contextual nuances. BART, a transformer-based model, performs exceptionally well in benchmark testing. The ROUGE score-based evaluation points out the algorithms' ability to generate coherent and relevant summaries. These findings illustrate the capability of advanced summarization techniques to aid information processing, making them important tools for researchers, professionals, and the public in need of effective access to important information.

Keywords - *Lexical Scoring, Semantic Similarity Measures, Machine Learning-Based Approaches, Rule-Based Approaches.*

1. Introduction

The runaway expansion of information in the current digital era, fueled by the World Wide Web (WWW), newspapers, emails, and databases, has resulted in an immense amount of textual information. Information overload has come from the significant difficulties in meaningfully understanding and using information [1]. This problem has long been addressed by document retrieval systems, which enable users to get relevant sections of documents using query searches. However, as retrieval techniques advance, the volume of data recovered often exceeds human capacity for analysis. A straightforward Google search query, for instance, may yield millions of results, requiring the development of processes for condensing and abstracting large volumes of

text [2]. Because of this, one of the main areas of study within the broader framework of natural language processing (NLP) is automated text summarization (ATS). Text summary aims to preserve the main ideas of long publications while condensing them into brief, relevant summaries. It is accomplished by using semantic or context-based processing methods to extract salient characteristics from documents and present them in a legible and consistent way [3]. When readers need to gain quick insights without reading lengthy text, this function is quite useful. The two main categories of summarization techniques are extractive and abstractive procedures. While abstractive summary creates new sentences to describe the material, extractive summarization takes important phrases or sentences from the original text and gathers them. Both strategies have advantages and disadvantages, and the choice of method is contingent upon the particular demands of the assignment [4].

Text summarization is centered on natural language processing (NLP), a multidisciplinary discipline that lies between computer science and artificial intelligence. The goal of natural language processing (NLP) is to enable computers to interact with, understand, and generate human language so that they can carry out activities like sentiment analysis, machine translation, and summarization [5]. Given the ambiguity and complexity of human language, natural language processing (NLP) employs sophisticated techniques such as statistical models, machine learning, deep learning, and rule-based procedures to extract important information and provide natural language output [6].

A number of methods, such as TextRank, Seq2Seq, and BART, have been suggested for text summarization. A graph-based method called TextRank uses connectedness to identify important components in text by modeling it as a word and sentence graph [7]. Seq2Seq, a deep learning algorithm, employs an encoder-decoder model to generate high-quality summaries. BART, a transformer model, reads text bidirectionally and attains state-of-the-art performance on benchmark datasets [8]. These algorithms, quantified in terms of metrics like ROUGE scores, demonstrate the ability of ATS systems to alleviate information overload, boost productivity, and provide actionable insights in applications like information retrieval, recommendation systems, and knowledge management. By enabling users to navigate large volumes of textual data quickly, text summarization continues to be a valuable tool in the big data era [9].

The research organized in Section 1 which is introduction of studies and Section 2 represents literature work. There are discussed about data and method with architecture of proposed system in Section 3 and result analysis in Section 4. With conclusion of research discussed in Section 5.

2. Literature Review

Text summarization has been a point of interest in natural language processing (NLP), with earlier work being mostly extractive in nature compared to abstractive. Extractive summarization, which involves the extraction and concatenation of important sentences from the original text, has been a widely researched area due to ease of use and effectiveness. One of the earliest contributions to

this area was by Luhn [10], who hypothesized that word frequency of occurrence in a document is directly proportional to the significance of its ideas. He outlined a process for scoring sentences by word frequency, choosing the highest-scoring sentences for the summary. On this basis, Edmunson proposed adding location-based, title-based, and cue-word-based features to determine key sentences, based on the hypothesis that key information is mostly in the front of a document. These statistical approaches, however, do not consider semantic sentence relationships, thus restricting them to extracting contextual meaning that is not profound [11].

To address this limitation, introduced a query-based summarization technique, where sentences are extracted based on their relevance to a user-provided query. This method assigns higher importance to sentences containing words that overlap with the query, ensuring the summary aligns with user intent. Goldstein [12] further combined statistical and linguistic features to evaluate and summarize news articles, enhancing the quality of extracted summaries. Another approach involves clustering sentences based on semantic similarity before applying extraction techniques. Proposed techniques using the K-means algorithm to group sentences, selecting representative ones from each cluster to form the summary. This method leverages semantic distance to ensure diversity and coverage in the summary [13].

Lexical chains, introduced another significant contribution to summarization. Lexical chains connect related words together based on their semantic relationships to highlight key ideas in the text. This was furthered by creating a linear-time lexical chain creation method that addresses issues like anaphora resolution and proper nouns. By building meta-chains of noun senses in WordNet, their technique makes it possible to retrieve crucial information quickly [14].

In summarization, graph-based techniques have also become more and more common. Semantic networks are created by the subject-object-predicate (SOP) triples, which capture the connections between key ideas in the text. Their approach provides a strong substitute for conventional techniques by identifying and extracting significant information from these networks. Pushpak Bhattacharyya [15] is comparable. Presented a WordNet-based summarization technique that uses WordNet synsets to build and weight a sub-graph of the document. This approach creates succinct summaries by utilizing language linkages.

In summary, text summarizing methods have advanced significantly, utilizing linguistic and statistical methods to tackle information extraction and compression challenges. While early methods used word frequency and sentence location, more recent approaches used lexical chains, graph models, and semantic analysis to provide more accurate and contextually rich summaries. [16]

3. Methodology

DATASET

The CNN/Daily Mail dataset is a huge corpus of news articles and synopses used in text summarization tasks. The dataset comprehends over 300,000 article-summary pairs, and hence it

is a appreciated source for training and testing summarization methods. The dataset is miscellaneous in subject matter, with politics, entertainment, sports, and other subjects. Articles and summaries are principally from the CNN and Daily Mail news websites, with further sources [17]. The summaries in the dataset are jobwise written by reporters and are typically shorter than the unique articles. This aspect makes the dataset perfect for training summarization approaches. The CNN/Daily Mail dataset has been a foremost force in studies in the turf of text summarization concluded the years and has been widely employed in research studies.

DATA PREPROCESSING

Preprocessing is essential while processing the CNN/Daily Mail dataset. The area is to preprocess and transform raw data to a appropriate format in order to evaluate or model. The most critical preprocessing tasks used in this dataset are recorded below:

Data Cleaning

This procedure entails the exclusion of irrelevant or duplicate data and handling missing or imperfect data. Figure 1 shows major features of the dataset afterward cleaning. Some of the most widely used techniques are:

- Eliminating stop words (e.g., "the," "and").
- Stemming and lemmatizing text to lower words to their base form.
- Eliminating unwanted characters, symbols, or formatting.

	text	y
0	LONDON, England (Reuters) – Harry Potter star...	Harry Potter star Daniel Radcliffe gets £20M f...
1	Editor's note: In our Behind the Scenes series...	Mentally ill inmates in Miami are housed on th...
2	MINNEAPOLIS, Minnesota (CNN) – Drivers who we...	NEW: "I thought I was going to die," driver sa...
3	WASHINGTON (CNN) – Doctors removed five small...	Five small polyps found during procedure; "non...
4	(CNN) – The National Football League has ind...	NEW: NFL chief, Atlanta Falcons owner critical...

Figure 1 Shown Dataset

Tokenization

Tokenization is the process of dividing text into smaller pieces, e.g., words or sentences. This is usually done using tokenizers that can handle punctuation, numbers, and special characters. Tokenization is needed for text preparation for further analysis.

Text Normalization

The process of converting text into a uniform and standard format to guarantee consistency and facilitate analysis is known as text normalization. In many applications using natural language

processing (NLP), this job is essential. The dataset's word frequency distribution following normalizing is shown in Figure 2. Text normalization sometimes entails transforming the text to lowercase, eliminating accents and other diacritical marks, and enlarging contractions, such "don't" to "do not." By standardizing the language and removing superfluous variants, these procedures increase the efficacy of subsequent processing and analysis.

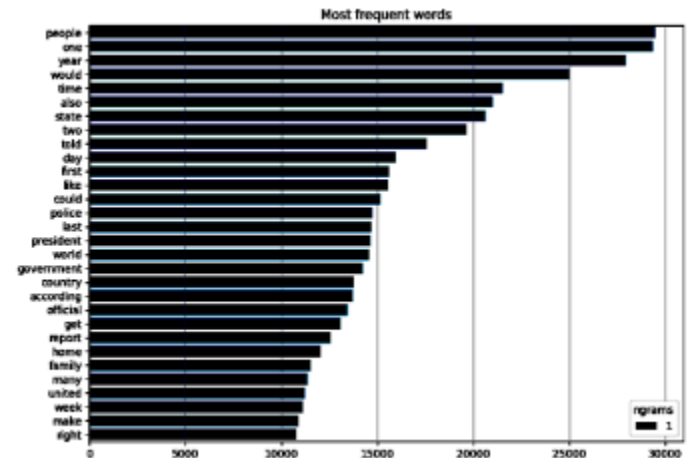


Figure 2 Dataset words frequency representation

POS Tagging

Part-of-speech (POS) tagging includes grammatical tags (e.g., noun, verb, adjective) to every word of the text. POS tagging is useful in syntactic analysis of the text and is useful for downstream tasks such as sentiment analysis and named entity recognition. Figure 3 shows the analysis of the text length distribution of the dataset after preprocessing.

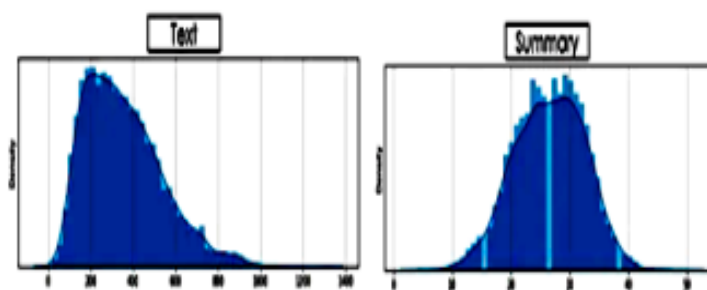


Figure 3 Text length distribution

Types of Summarization Techniques

Extraction-Based Summarization

Extraction-based summarization is the selection and extraction of significant expressions or sentences from the input text to form a summary. One of its strengths is that it may be able to maintain the original wording and meaning of the text, which can be extremely significant in

certain cases. The process is highly appropriate for summarizing factual information, such as news reports, and is relatively easy to implement. There are various methods of implementing extraction-based summarization, such as frequency-based and machine learning-based methods: These methods select significant sentences based on statistical criteria such as the frequency of words or the length of sentences. These are trained algorithms that select significant sentences based on criteria such as sentence length, position in the text, and the presence of significant words. Although extraction-based summarization may sometimes generate summaries that are unstructured and incoherent in the sense that the selected sentences may not be well-related to each other, it may also struggle with complex texts that must be interpreted and understood at a deeper level.

Abstraction-Based Summarization

Abstraction-based summarization generates a summary that is not necessarily in the same language as the original text but instead summarizes its main ideas and concepts in a more general way. This is achieved through natural language generation techniques that generate new sentences that express the main ideas more concisely and clearly.

This is most suitable for technical or scientific texts because it summarizes complex ideas and idea relationships well. Abstraction-based summarization, however, has some drawbacks:

- It requires more training data and computational resources than extraction-based summarization.
- The quality of the natural language generation techniques employed can significantly impact the quality of the summary, and hence optimization is challenging.

Algorithm Selection

TextRank Model: Graph-based approach based on PageRank for extracting most relevant sentences after preprocessing (stop word removal, stemming, lowercasing).

TextRank(input_text):

input_text = Preprocess(input_text):

Remove stop words, apply stemming, lowercase

sentence_graph = Build_Sentence_Graph(input_text)

Compute_Similarity(sentence_graph)

ranked_sentences = Apply_PageRank(sentence_graph)

summary = Extract_Top_Sentences(ranked_sentences)

RETURN summary

Seq2Seq Model: Encoder-decoder architecture-based approach. Preprocessing (conversion and tokenization) and model training to predict the summary are used. ROUGE and BLEU scores are

used for evaluation.

1. *Preprocess input and target text by tokenizing and converting to numerical format.*
2. *Build a Seq2Seq model with an encoder-decoder architecture and attention mechanism.*
3. *Train the model on the preprocessed input and target text data.*
4. *Evaluate the model performance using metrics like ROUGE and BLEU.*
5. *Return the performance evaluation results to assess summarization quality.*

BART Model: Encoder-decoder architecture-based, trained on a summarization dataset and tested using ROUGE and BLEU scores.

1. *Preprocess input and target text by tokenizing and converting to numerical format.*
2. *Build a BART model with an encoder-decoder architecture.*
3. *Fine-tune the BART model on the preprocessed input and target text data.*
4. *Evaluate the model performance using metrics like ROUGE and BLEU.*
5. *Return the performance evaluation results to assess summarization quality.*

TextRank is a graph ranking algorithm used for text summarization. It models text as a graph in which each node is a word or a phrase and the edges between them signify the similarity relations. PageRank algorithm is then used to rank the nodes, in which the edge weights are calculated based on the semantic similarity of the nodes. Preprocessing is performed on the input text before making a graph. Preprocessing involves stop word removal, punctuation, and unwanted noise. After preprocessing the text, a graph is constructed with sentences or words as nodes and edges formed based on similarity measurements like Cosine similarity, Jaccard similarity, or other semantic distance measurements.

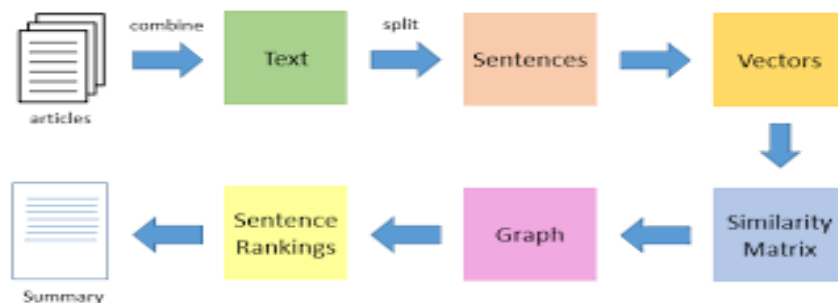


Figure 4 Working procedure of Textrank

After building the graph, the PageRank method, which is modified in the case of TextRank for nodes that are similar to their neighbors, gives ranking scores to each node. The ranking of a node

is calculated by taking the weighted average of the ranks of its nearby nodes. How similar they are to one another determines the weights. The summary is created by selecting the nodes that score highest in the list and represent the most significant words or phrases. A succinct and insightful synopsis that accurately captures the essence of the original text is then created by combining these key terms. Fig. 4's flow diagram provides a visual representation of the TextRank process. The purpose of Seq2Seq models is to create a mapping among input and output data sequences, including the sequence of words in a source language and its equivalent order in the target language. The encoder and the decoder are the two primary parts of these models. The given input sequencing is processed by the encoder and supplied to the decoder via a fixed-length vector representation. To build the target sequence token by token, the decoder uses this encrypted representation in conjunction with outputs that have already been generated. Recurrent neural networks (RNNs), such as Gated Recurrent Units (GRU) or Long Short-Term Memory (LSTM), are commonly employed as encoders.

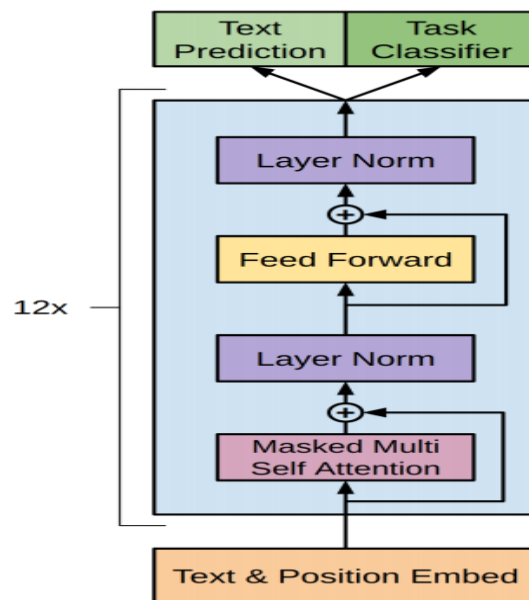


Figure 5 Architecture of BART

The encoder's ultimate hidden condition serves as a summarized depiction of the whole input sequence, as well as these connections process the input sequence a single token at a period of time updating their concealed states at each step. The decoder, which is structurally distinct, is also an RNN that produces output tokens repeatedly, using the encrypted representation from the encoder as well as the earlier produced token as input. The decoder keeps producing tokens until it encounters an end-of-sequence token or reaches a predetermined maximum length. The transformer design, which models interactions between items in a sequence, such as words in a

phrase, using attention processes, is expanded upon by BART (Bidirectional and Auto-Regressive Transformer). In contrast to conventional Seq2Seq models, BART combines an auto-regressive decoder with a bidirectional encoder. Similar to Seq2Seq, the encoder converts input sequences into a fixed-length vector format. By processing input both forward and backward, BART's encoder outperforms traditional Seq2Seq encoders in natural language processing (NLP) tasks by more accurately capturing complex relationships between tokens. The BART decoder functions similarly to a Seq2Seq decoder, although using an auto-regressive technique and generating each output token based on the previously formed tokens. As a result, the model can identify connections between output tokens, resulting in more coherent and contextually accurate language. BART is pre-trained using a denoising autoencoder approach, which randomly masks or rearranges parts of the input text, before training the model to reconstruct the original content.

. As a result of this pre-training, BART develops a deeper understanding of real language, making it incredibly effective for a range of NLP tasks. Fig. 5 shows the BART flow diagram.

Proposed Method

There are several automatic procedures involved in turning an HTML document into plain text and summarizing its contents. A succinct and excellent summary may be produced by using an HTML-to-text parser with pre-processing, sentence score, and normalization algorithms.

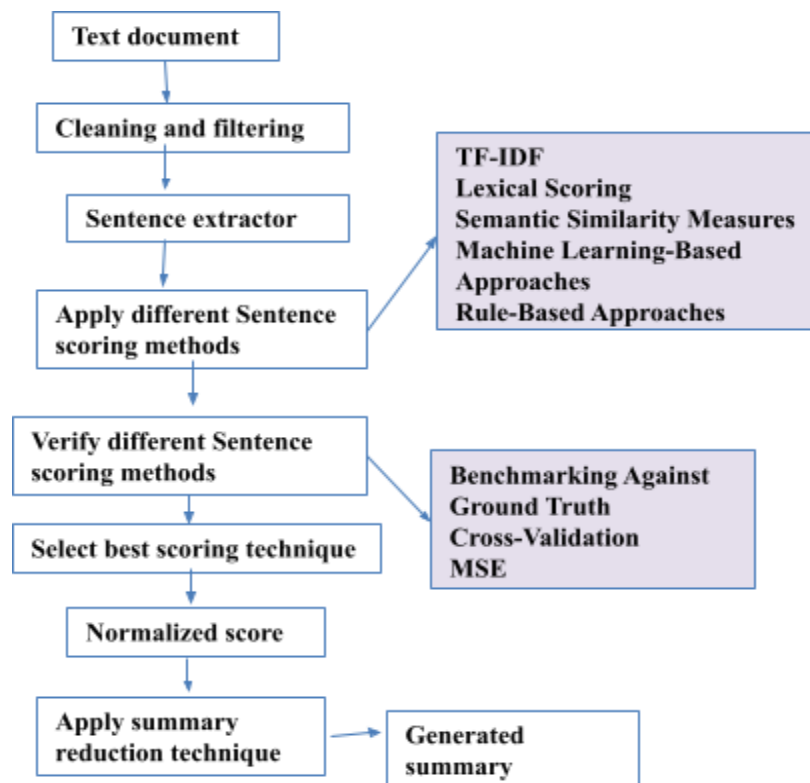


Figure 6 DFD of the Proposed System

In figure 6 explain a systematized pipeline for a text summarizing or NLP-based valuation system is showed in the diagram. It starts with basic data processing, in which preprocessing procedures including tokenization, stopword removal, and normalization are functional to the incoming text. Then, feature extraction approaches are used, such as rule-based methods, lexical scoring, semantic similarity measures, machine learning-based procedures, and TF-IDF. The summarization model accepts these extracted appearances as input and uses them to evaluate the data and produce summaries. The fashioned summaries are then assessed using benchmarking procedures against ground truth data to guarantee quality and accuracy. Performance measures like Mean Squared Error (MSE) and cross-validation methods aid in evaluating the model's efficacy. Eventually, the system generates a well-optimized and cohesive summary following thorough testing and enhancement. For increased performance in NLP applications, this technique guarantees that the summarizing process is data-driven, methodically measured, and continually improved.

4. Result Analysis

In table1 set of sentences (5, 10, 15, 20, 25) were reserved from corpus. Then their sentence scores calculated by applying numerous scoring procedures as declared in the table. Yi grants the average actual sentence score of (5, 10, 15, 20, 25) sentences correspondingly.

Table 1: On behalf of Rule based approach through Lexical approach as well as TF-IDF

No. of sentences	Rule based approach			Lexical approach			TF-IDF		
	Yi (actual value of observation)	Y'i (predicted value of observation)	(Yi - Y'i)	Yi (actual value of observation)	Y'i (predicted value of observation)	(Yi - Y'i)	Yi (actual value of observation)	Y'i (predicted value of observation)	(Yi - Y'i)
5	23	30	7	23	32	9	24	33	8
10	26	33	7	27	37	10	28	39	11
15	25	28	3	27	31	4	27	30	3
20	29	32	4	29	32	5	29	35	6
25	30	33	4	30	33	4	30	33	3
		MSE	26.4		MSE	47.6		MSE	47.8

Below is a summary of the findings utilizing the TextRank, BART, and Seq2Seq algorithms.

rouge1: 0.21 | rouge2: 0.07 | rougeL: 0.07 --> avg rouge: 0.15

Figure 7 TextRank Score

rouge1: 0.08 | rouge2: 0.01 | rougeL: 0.01 --> avg rouge: 0.05

Figure 8 Seq2Seq Score

rouge1: 0.3 | rouge2: 0.11 | rougeL: 0.11 --> avg rouge: 0.21

Figure 9 BERT Score

5. Conclusion and Future Work

A key area in NLP is text summarizing, which purposes to provide succinct and logical summaries of widespread texts. Both abstractive and extractive approaches can be used. TextRank, Seq2Seq, and BART are samples of recent developments in deep learning models that have importantly enhanced summarization enactment. The PageRank technique is used by the graph-based algorithm TextRank to find important phrases and provide extractive summaries. Conversely, Seq2Seq models use neural networks to train on document-summary pairings in order to produce abstractive summaries. Text summarization is one of the many NLP applications where BART, a pre-trained transformer model, has shown cutting-edge performance. Every one of these approaches has unique benefits and drawbacks. Although TextRank is easy to use and doesn't require training data, it could have trouble creating fluid summaries. Although Seq2Seq models require a large amount of training data and can yield generic results, they give more organic and cohesive abstractive summaries. Although BART is very effective and versatile in producing coherent and abstractive summaries, it is computationally costly and necessitates extensive dataset pre-training. By studying the various sentence, scoring techniques it comes in the conclusion that Rule based approach gives accurate sentence score as compared to other techniques. Then further, by applying appropriate normalization technique a more accurate summary is generated. There is a lot of promise for future text summarizing research. One interesting approach is multi-document summarization, which entails analyzing several texts on the same subject. Another possible avenue is the creation of domain-specific summarization models suited to particular disciplines. The user experience may be improved via interactive summarizing, in which summaries are improved for more accuracy and customisation based on user input. Furthermore, developing multilingual summarization models and enhancing assessment measures continue to be important research topics.

References -

- [1] Automatic Text Summarization and Keyword Extraction using Natural Language Processing, "Avinash Payak; Saurabh Rai; Kanishka Shrivastava; Reshma Gulwani", IEEE, 02-04 July 2021, 10.1109/ICESC48915.2020.9155852.
- [2] James Allan, Jaime Carbonell, George Doddington, "Domain Specific Document Summarization by sentence extraction", Journal of King Saud University - Computer and Information Sciences, Volume 32, Issue 10, December 2021, Pages 1227-1228.
- [3] Suad Alhojely, "A scalable summarization system using robust NLP", 2020 International Conference on Computational Science and Computational Intelligence (CSCI), 978-1-7281-7624-6/20/\$31.00 ©2020 IEEE.
- [4] Breck Baldwin and Thomas S. Morton, "Multi-Document Summarization using Sentence Extraction for user query", Appl. Sci. 2020, 12(9), 4479.
- [5] Single Document Text Summarization Algorithm using semantic similarity, International Journal of Computer Applications (0975 – 8887) Volume 17– No.2, March 2020.
- [6] A. Siddharthan, A. Nenkova, and K. McKeown. Syntactic simplification for improving
- [7] Dhanda, Namrata, and Kapil Kumar Gupta. "A Novel Approach to Text Summarization Using Machine Learning." *Asian Journal of Research in Computer Science* 17, no. 4 (2024): 95-104.
- [8] Godoy, Wagner F., José A. Fabri, Rodrigo HC Palácios, Márcio Mendonça, Janaína FS Gonçalves, and Luiz OM Moraes. "Improving Texts into Powerful Communication Tools with Classic Techniques and Deep Learning."
- [9] Sharma, Pallavi, and Min Chen. "Text Summarization and Keyword Extraction." In *2023 14th IIAI International Congress on Advanced Applied Informatics (IIAI-AAI)*, pp. 369-372. IEEE, 2023.
- [10] Sood, S., Ukrit, M.F. and Parasar, R., 2023, May. Analysis of extractive text summarisation techniques for multilingual texts. In *2023 International Conference on Advances in Computing, Communication and Applied Informatics (ACCAI)* (pp. 1-8). IEEE.
- [11] Balaneji, Farshid. "Language as a Lens: A Hybrid Text Summarization and Sentiment Analysis Approach for Multiclass Stock Return Prediction." In *Intelligent Systems Conference*, pp. 429-448. Cham: Springer Nature Switzerland, 2024.
- [12] Ingole, Akhilesh, Prathamesh Khude, Sanket Kittad, Vishakha Parmar, and Archana Ghotkar. "Strategic Brand Sentiment Analysis for Competitive Surveillance." In *2024 3rd International Conference on Applied Artificial Intelligence and Computing (ICAAIC)*, pp. 1114-1119. IEEE, 2024.
- [13] Ahmad, Ayaan, and Dr Ng Kok Why. "Automated Grading Using Natural Language Processing and Semantic Analysis." Available at SSRN 4999531.
- [14] Deekshan, S., and Arjun Dev PK. "Detection and summarization of honest reviews using text

mining." In 2022 8th International Conference on Smart Structures and Systems (ICSSS), pp. 01-05. IEEE, 2022.

- [15] Zhang, Weijia, Jia-Hong Huang, Svitlana Vakulenko, Yumo Xu, Thilina Rajapakse, and Evangelos Kanoulas. "Beyond relevant documents: A knowledge-intensive approach for query-focused summarization using large language models." In International Conference on Pattern Recognition, pp. 89-104. Springer, Cham, 2025.
- [16] Wu, Chaoyi, Pengcheng Qiu, Jinxin Liu, Hongfei Gu, Na Li, Ya Zhang, Yanfeng Wang, and Weidi Xie. "Towards evaluating and building versatile large language models for medicine." npj Digital Medicine 8, no. 1 (2025): 58.
- [17] Van Veen, Dave, Cara Van Uden, Louis Blankemeier, Jean-Benoit Delbrouck, Asad Aali, Christian Bluethgen, Anuj Pareek et al. "Adapted large language models can outperform medical experts in clinical text summarization." Nature medicine 30, no. 4 (2024): 1134-1142.