

Predictive Text Summarization: Improving Efficiency and Effectiveness

Rehan Sajid¹, *Muhammad Munwar Iqbal², Shabana Ramzan³

^{1,2} Department of Computer Science, University of Engineering and Technology, Taxila

³Department Computer Science & IT, Govt Sadiq College Women University Bahawalpur.

ABSTRACT

Automatic text summarization is essential for knowledge extraction and text classification. It provides the solution to the issue of data and information overloading. The volume of text data accessible has risen significantly in recent years from several sources. A large amount of text contains a variety of detail and insights that must be effectively documented. Text Summarization is the act of shrinking a text while maintaining the information values and transforming them into concise summaries that state the document's key objective. The proposed model extraction-driven text summarization entails extracting high-ranking sentences from a document based on term and phrase attributes and combining them to provide a description. The Fuzzy inference engine is used to model the document's summary and calculated by the relative value of the sentences in the document. The semantic approach to utilizing Latent Semantic Processing is explored, and the Fuzzy logic Extraction approach for text summarization. The proposed system suggested system has an average recall of 44.51, an average precision of 90.83, and an f-measure of 67.66. Our proposed summarizer's overall enhanced recall, precision, and f-measure efficiency than the fuzzy-based summarizer.

Keywords: Fuzzy Inferences, Semantic Analysis, Document Classification, Automatic Text Summarization

Author's Contribution

¹ Data analysis, interpretation and manuscript writing, Active participation in data collection. ²Conception, synthesis, planning of research, ³ Interpretation and discussion

Address of Correspondence

Muhammad Munwar Iqbal
Email: munwar.iq@uettaxila.edu.pk

Article info.

Received: January 26, 2021
Accepted: June 13, 2021
Published: June 30, 2021

Cite this article: Sajid R, Iqbal MM, Ramazan S. Predictive Text Summarization: Improving Efficiency and Effectiveness. *J. inf. commun. technol. robot. appl.*2021; 12(1):11-19.

Funding Source: Nil

Conflict of Interest: Nil

INTRODUCTION

The amount of information available on the Internet has increased tremendously over time. There is no organized presence of all the available data. Desired information can be gotten from unorganized data, and an enormous amount of time and resources are needed. The process in which the original text generates relevant and meaningful information is called automatic text summarization (ATS). The main goal of the text summary is to automatically

extract useful information from the bulk of meaningless data, rather than reading it all and manually summarizing it, which takes a lot of time and resources. ATS research began in the mid-90s, but ATS is still considered the major and challenging task in the NLP and AI field.

The three methods for text summarization are extractive, abstractive, and hybrid. Extracting the same but significant sentences from the given text or document is

an extractive text summary because it does not produce new phrases and does not combine phrases to form a new one, which is its limitation. The steps involved in extractive text summarization to get the target summary are preprocessing the source document and post-processing. Instead of extracting exact phrases from the input text, abstractive text summarization generates new meaningful phrases by using different methods to understand the text's main idea. Two or more phrases are combined, and new words are added to form a new phrase, but human-like summaries are produced as output. Getting the desired summary also involves preprocessing and post-processing of the input document. Hybrid summarization is a combination of extractive and abstractive summarization, in which the steps involved are preprocessing, sentence extraction, summary generation, and post-processing.

Interpretation, transformation, and generation are the steps involved in text summarization. The interpretation phase creates adequate semantic representations using natural language processing (NLP) from the original text document. The second step is synthesis, which takes as an input the results of the interpretation step and extracts sufficient information to provide an adequate overview. Generation is the last stage in the process of text summary, producing the Summary as an output with the generation of natural language (NLG).

After implementing all these steps, a brief and summarized text will be produced as an output with less repetition of the same words. Different approaches have been used to build an automated text summarization method based on the genesis, space, and NLTK. Firstly, text representation models represent the input document for processing purposes. The vector, matrix, n-gram, and topic model may be several text representations models. Secondly, linguistic analysis and processing methods such as preprocessing, parsing, grammar, disclosure analysis, the resemblance of sentences, and natural language generation may be used in various ATS levels. Finally, soft computation methods such as machine learning algorithms, optimization algorithms, and fuzzy logic may be used for applying text summarization. Various evaluation methods may automatically evaluate summaries like accuracy, recall, F-measure, ROUGE-1, ROUGE-2, or ROUGE-L metrics.

Latent Semantic Analysis (LSA)

LSA is a completely automated mathematical/statistical technique for extracting and inferring predicted semantic use relationships of terms in discourse passages. It is not a conventional natural language processing or artificial intelligence program; it does not use any human-created dictionaries, information bases, semantic networks, grammars, syntactic parsers, or morphologies and only recognizes raw text parsed into terms identified as specific character strings and divided into coherent passages or examples like sentences or paragraphs as input. The first move is to visualize the document as a matrix, with each row representing a single word and each column representing a text passage or other meaning. The frequency at which the term in its row occurs in the passage denoted by its column is registered in each cell. The cell entries are then subjected to a preliminary transformation, the specifics of which will be described later, in which each cell frequency is weighted by a mechanism that expresses both the value of the terms in the specific passage and the degree to which the world type carries knowledge in the discourse domain in general. The matrix is then subjected to singular value decomposition (SVD) by LSA. It is a form of factor analysis or, to be more specific, a statistical generalization to which factor analysis is a subset. An SVD decomposition decomposes a rectangular matrix into the sum of three other matrices. When the three elements are matrix-multiplied, the initial matrix is restored. One dimension matrix represents the original row entities as vectors of derived orthogonal factor values, another matrix defines the original column entities in the same manner, and the third is a diagonal matrix comprising scaling values. There is statistical proof that every matrix can be completely decomposed using no more variables than the initial matrix's smallest dimension. The restored matrix is a least-squares best fit where fewer variables are included than are needed. Delete coefficients in the diagonal matrix to reduce the dimensionality and start from the initial value.

This study aims to improve document and sentence representation to improve sentence extraction for single word extractive summarization. Improved text representations aid in the development of better extract summaries. Sentence raking is an essential part of the

extractive summarization process. We will enhance the sentence rating process by improving the sentence representation. Sentence rating assists in the production of better sentence scoring for better extract summaries. It maximizes sentence collection. Repetitive sentences must be identified, which improves the consistency of multi-document summaries. We concentrate on using semantics and attention-based compositional features latent in the text to enhance text representation.

BACKGROUND

Established approaches to summarization are commonly either extractive or abstractive. The model extracts passages from the input text in extractive summarization and blends them to form a shorter summary, often with a post-processing stage to ensure the final coherence of the output [1, 2]. Although extractive models are typically resilient and generate coherent summaries, they are unable to construct succinct summaries that use new phrases to paraphrase the source text.

Abstractive summarization helps the model, not in the source document, to paraphrase the source document to construct succinct summaries of sentences. The state-of-the-art abstractive summary models focus attention on sequence-to-sequence models [3, 4]. Extensions to this paradigm include a mechanism of self-attention or textual engagement through reinforcement learning; and generative adversarial networks to produce more natural summaries [5].

In this article, the author [6] explains a Bayesian summarizer for biomedical text documents. Initially, the Bayesian summarizer maps the input text to the Universal Medical Language System (UMLS) and then chooses the relevant ones to be used as classification functions [7]. It defines the most relevant concepts of the text, and to choose the most informative material according to the distribution of these concepts, the author incorporates numerous feature selection approaches [8]. This research reveals that the Bayesian biomedical summarizer can boost the efficiency of summarization by using an effective feature selection method. The scientist conducts detailed analyses on a corpus of research articles in the biomedical area. The results reveal that, based on the Recall-Oriented Understudy for Gisting Evaluation (ROUGE) metrics, the Bayesian summarizer outperforms

the biomedical summarizers that focus on the frequency of terms, domain-independent, and baseline approaches [9]. Besides, the findings indicate that using the measure of meaningfulness and considering the associations of concepts in the selection process of the function leads to a substantial improvement in the summary output.

Bayesian biomedical summarizer is limited, yet more research work is required to deal with multi-document and query-focused summarization [7]. A more complex form of redundancy reduction should be explored to do so. It seems like in the sense of concept-based summarization. There is far more space for learning the Helmholtz theorem from the Gestalt theory. The author [10] also modeled document sentences using the Helmholtz theory as a small-world network and studied implementations such as text summarization. For concept-based biomedical text summarization, future work can include exploring this modeling method.

2.1. Extraction Methodologies

The bulk of text summarization literature is devoted to extractive summarization methods. In such extractive structures, sentence extraction is a critical stage. The aim is to find a subset of sentences representing the entire set and containing all of the information. Orthodox extractive summarization methods classify sentences using manually constructed features such as sentence position and length [11-13], title terms, the inclusion of proper nouns, material features such as phrase frequency [14], and occurrence features such as action nouns [15]. In general, sentences are given a saliency score that indicates the degree to which the features mentioned above are present. Kupiec et al. [16] pick description-worthy sentences using binary classifiers. Hidden Markov Models were investigated by Conroy and O'Leary [17], and graph-based algorithms for choosing salient sentences were implemented [18]. Regardless of the sentence ranking models, the efficiency of the summarization scheme using the sentence ranking method is heavily influenced by function engineering [19, 20]. Document-dependent features are location and term frequency. Document-independent features are length, stop-word ratio, and word polarity. The fact that a statement can be called Summary worthy regardless of which text it appears in it. It is also revealed by document-independent functionality.

2.1.1 The TF-IDF technique (Term Frequency-Inverse Document Frequency)

The number of times a phrase appears in a sentence is referred to as its "sentence frequency." Then, based on the query's closeness to the sentence vectors, the most relevant phrases are selected for inclusion in summary.

2.1.2 Based on clusters

Documents are often written in an orderly fashion, with each section addressing a separate subject. They are usually divided into parts, either expressly or implicitly. It's natural to assume that summaries should include a wide range of "themes" that exist across the summaries' sources. The clustering method used some summaries, including this feature. Clustering is virtually a must for summaries of large collections of documents because of the wide range of subjects covered. The selection of sentences is based on their resemblance to the cluster C_i 's topic. The position of the sentence in the text is the next consideration in sentence selection (L_i). For inclusion in summaries, it is closer to the beginning of a sentence. Lastly, a sentence's resemblance to the document's opening sentence is a criterion that raises its score (F_i). Three aspects contribute to a sentence's total score:

$$S_i = W_1 * C_i + W_2 * F_i + W_3 * L_i$$

2.1.3 It's a Graph-Based Approach

Thematic identification may be facilitated using graph-theoretic representations of passages. These nodes in the undirected network reflect sentences that have undergone standard preprocessing processes, such as stemming and stop word removal. Every sentence has a node. An edge connects two phrases if they include similar terms. Sentences may be picked from the relevant subgraphs for query-specific summaries, while for general summaries, sentences can be taken from all the subgraphs. There are a lot of relevant sentences in the partition, and those with a high cardinality receive the most attention when summarizing.

2.1.4 Strategy Based on Machine Learning

To model summarization, sentences are divided into summary and non-summary, depending on their characteristics. The Bayes' rule is used to learn the categorization probabilities statistically:

$$P(s \in S | C_1, C_2, \dots, C_N) = P(C_1, C_2, \dots, C_N | s \in S) * P(s \in S) / P(C_1, C_2, \dots, C_N)$$

The characteristics $C_1, C_2, \dots$, and C_N are utilized to classify the phrase S in the document collection. To construct a summary, S is the target, and $P(s \in S | C_1, C_2, \dots, C_N)$ is the probability of selecting a phrase with attributes $C_1, C_2, \dots, C_N$.

2.1.5 The LSA method

Since Singular Value Decomposition is used to document-word matrices that are semantically connected but do not share common terms, this technique is known as Latent Semantic Analysis. When used in a solitary context, words are often seen together take on a new meaning. Content sentences and topic words from texts may be extracted using this technique. Because conceptual (or semantic) relations in the human brain are naturally recorded in LSAs, rather than in word vectors, summarization may be done more easily.

2.1.6 Using Neural Networks to Summarize the Text

For this strategy to operate, you must first train neural networks to recognize the different sorts of phrases that belong in a synopsis. The reader is responsible for this. When a neural network analyses a text, it learns what patterns are there and what patterns are not present. Feedforward neural network, which has been demonstrated to be a universal function approximator, is used in this system.

2.1.7 Fuzzy Logic-Based Text Summarization

The fuzzy algorithm considers all a text's characteristics, such as its resemblance to little, its sentence length, and its proximity to a keyword. Then, the rules for summarizing are entered into the system's knowledge base. Each of the output sentences is given an integer value from zero to one determined by the properties of the sentences and the rules that exist in the knowledge base. Afterward, the output value decides how important a statement will be in the final summary.

2.1.8 Query-based text extraction and summary

Query-based text summarization systems rate sentences based on the frequency of keywords in each document. More points are awarded to sentences that include query phrases than to those that just contain a single query term. Afterward, the most frequently occurring sentences and their structural context are included in the summary. It is possible to extract parts or subsections of text. Such excerpts are combined to produce the final summary.

2.2. Abstractive Methodologies

An abstractive summary is generated using Structure-based Abstractive Summarization, which populates prominent sentences in a predefined structure without losing their meaning. Templates, tree-based structures, ontology-based structures, lead and body word structures, and rule-based structures are among the predefined structures used. The extracted data is populated into a prototype to produce the final description in the template-based process [21]. With the aid of a parser, the tree-based approach extracts identical sentences from the source text and populates them into a tree structure that fits the predicate-argument structure [22]. The ontology-based approach preprocesses the necessary keywords from the source text and then maps them as concepts and relations using a predefined ontology, translated to a meaningful description [23]. The lead and body phrase strategies operate by substituting or adding information-rich related phrases from the body, known as stimuli, into the lead sentences. If the body phrase has a greater syntactic resemblance to the lead phrase and the detail is richer than the lead phrase, the body phrase replaces the lead phrase [24]. The rule-based approach represents text summaries of laws and categories. This module is fed rules to generate the requisite meaningful candidates, from which the best candidate is chosen and passed to the summary generation module to generate the Summary using the generation pattern [25].

Abstractive summarization based on semantics includes feeding the document's semantic representation into a natural language generation module to generate the desired description. This technique employs three distinct methods: (a) a multimodal semantic model, (b) an information item-based method, and (c) a Semantic Graph-based method. With the aid of an ontology, a multimodal semantic model builds a semantic model by using definitions and finding relationships between them. The next step is to use information density metrics to define the most relevant principles. In the final stage [26], the description is created from these key concepts. The knowledge item-based approach first distinguishes informative objects by doing syntactic analysis of the text. Sentences are generated from these objects using a sentence generator that follows the subject-verb-object form. The generated sentences are then graded using the

Document Frequency score as per criterion. The overview [27] comprises the highest-ranking sentences from this list. Techniques for abstractly summarizing information Abstractive strategies for summarization may be divided into two basic categories: structured and semantic.

2.2.1 An Approach Based on a Specific Structure

Cognitive schemas such as templates, extraction rules, and other structures such as a tree, ontology, lead, and body phrase structure are used in the structured-based method to encode the most significant information from documents.

2.2.1.1 Using a tree structure

The text and contents of a document are represented using a dependency tree in this method. Summary material may be selected using many different techniques, such as a topic intersection approach or an alignment method that takes into account the similarities and differences between the sentences in a pair of processed phrases. A language generator or an algorithm may be used to generate the summary. This technique lacks a comprehensive model that would include an abstract representation for content selection, which is a restriction.

2.2.1.2 Using a template-based approach

A template is used to represent the whole of a document in this method. Text samples are analyzed against linguistic patterns or extraction criteria to determine which template slots they should be mapped to. These short passages of text provide as a hint as to what the rest of the document will be like. The Information Extraction systems populate the templates with relevant text samples. The resulting summary has a high degree of coherence since it depends on pertinent information detected by the Internet Explorer (IE) system. This strategy can be used only if the source papers already have summary sentences. The work of multi-document summarization will be impossible for it to undertake if information regarding similarities and differences between several papers is required.

2.2.1.3. Based on an Ontology

Many scholars have tried to enhance the summarization process by using ontologies (knowledge bases). The vast majority of online content is domain-related, meaning it addresses the same issue or occurrence again. Ontologies can better depict the knowledge structure of

each area. The advantage of this technique is that it uses fuzzy ontology to handle data that a basic domain ontology can't handle. This method is only applicable to Chinese news.

2.2.1.4 The Lead and Body Phrase Technique

The identical grammatical head chunks in the lead and body sentence phrases are inserted and substituted to rebuild the lead sentence. This technique has the potential advantage of identifying adjustments to a lead sentence that are semantically suitable. There are various drawbacks to this approach. To begin with, mistakes in parsing reduce the fullness of a sentence by introducing grammaticality and repetition. It does not include an abstract representation for content selection, which would be an important component of a comprehensive model.

2.2.1.5 Methods Based on Rules

With this technique, the summary papers are presented categorically with a list of qualities. A content selection module employs information extraction criteria to determine the best available candidate to address one or more aspects of a category. Finally, summaries are constructed utilizing generating patterns for sentences. The strong aspect of this technique is that it has a possibility for developing summaries with better information density than the existing state-of-the-art. The biggest problem with this process is that all the rules and patterns are manually developed, which is laborious and time-consuming.

2.2.2 An Approach Based on Semantics

The semantic representation of a text is fed into a natural language generation (NLG) system in the semantic-based technique. This technique uses linguistic data to identify noun phrases and verb phrases.

2.2.2.1 Semantic model with several modes of expression

Multimodal texts are represented using a semantic model that contains ideas and relationships. The most significant ideas are ranked based on various criteria, and the chosen concepts are summarized in sentences. Abstract summaries produced by this framework have high coverage since they incorporate key textual and graphical elements from throughout the document. This framework's drawback is that it is reviewed manually by people. The framework should be evaluated automatically.

2.2.2.2 Method based on data items

Rather than sentences from the source documents, an abstract representation of source materials is used to produce summary content. Information Item, the text's smallest unit of cohesive information, serves as the abstract representation. This method's key advantage is that it produces concise summaries rich in information and minimizes redundancy. There are various drawbacks to using this method. Creating meaningful and grammatically correct phrases with several potential information pieces is tough. Second, poor parsing results in poor linguistic quality in summaries.

2.2.2.3 Graph-Based Semantic Analysis

Building a semantic graph (RSG) for the original content and then reducing the resultant semantic graph to a smaller size. From the reduced semantic network, it generates the final abstractive summary. The method's main advantage is that it creates grammatically accurate, succinct, and cohesive sentences. However, this approach can only summarize one document at a time.

RESEARCH METHODOLOGY

Using latent semantic analysis and a fuzzy logic system to extract sentence features, our proposed approach will increase the consistency of the Summary as shown in figure 1.

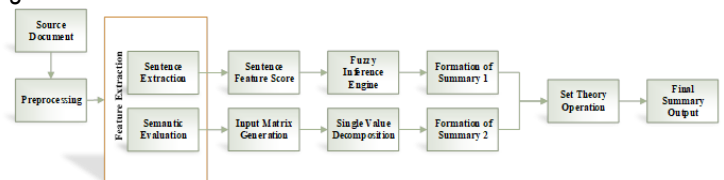


Figure 1: Proposed methodology of automatic text summarization

The machine reads the source text. The source document is fed into the proposed system after preprocessing. The feature extraction module is executed and extracted the sentence extraction. Simultaneously, semantics are also extracted in from the document features. The most representative sentences for a term are those with the highest index meaning. The sentences feature score is used for the fuzzy inference engine.

On the other hand, input Matix is generated from the semantic features. The highest-scoring sentences are

extracted as text summaries depending on the compression rate. Single valued decomposition is executed for the input generated matrix to develop the summary. A fuzzy inference engine is executed on the sentence feature extracted to formalize the text summary. We intersect all summaries and derive a collection of normal and uncommon sentences. The number of sentences in the Summary is calculated based on the compression rate and removes the uncommon sentences from the uncommon package. We integrated the semantic contents into the sentence function extraction Fuzzy summarization process. As a result, our suggested approach improves the Summary.

RESULT AND DISCUSSION

Our summaries have a high recall and accuracy relevance measure for manual assessment outcomes, which corresponds well with human judgment. As an assessment of experiment outcomes, we prefer recall and precision. We used a total of nine separate datasets shown in Table 1.

Table 1. Recall, Precision, and f-measure Values of fuzzy-based Summary.

Datasets	Recall	Precision	F-measure
	41.43	87.88	64.66
	42.76	92.30	67.30
	47.70	91.18	69.44
	33.33	82.58	58.04
	44.44	92.24	69.84
	46.03	87.88	67.06
	45.75	87.10	66.43
	43.10	84.85	64.07
	41.31	79.20	60.24
Average	42.87	87.24	65.23

The average fuzzy dependent summary recall is 42.87, the average accuracy is 87.24, and the average f-measure is 65.23. Precision, recall, and f-measure have their graphs and are shown in Figure 2.

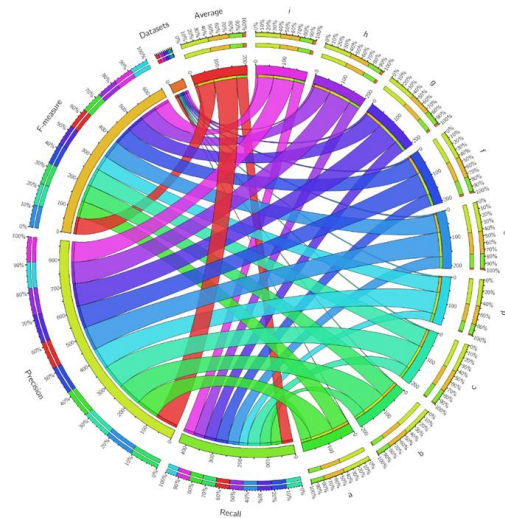


Figure 2. Assessment graph of Fuzzy based Summary

The recall, precision, and F-measure of 09 datasets that we used for fuzzy-based Summary are seen in the table. Recall, precision, and f-measure graphs are seen separately in Table 2.

Table 2. Recall, Precision, and f-measure Values of Proposed Summary

Datasets	Recall	Precision	F-measure
1	42.86	91.01	66.90
2	42.86	92.31	67.60
3	49.23	94.12	71.70
4	34.80	86.21	60.47
5	44.44	95.24	69.84
6	50.80	97.07	73.90
7	47.46	90.32	68.90
8	44.62	87.88	66.25
9	43.50	83.33	63.41
Average	44.51	90.83	67.66

The suggested summary's average recall is 44.51, its average precision is 90.83, and its average f-measure is 67.66. Precision, recall, and f-measure have their graphs and are shown in Figure 3.

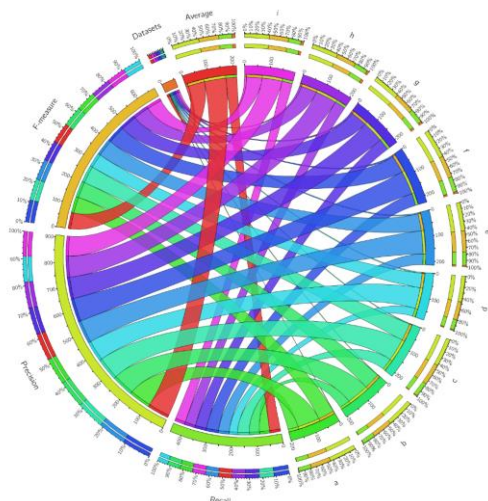


Figure 3. Assessment graph of Proposed Summary

CONCLUSION

Automatic summarization is a multi-step process that includes many sub-tasks. Each sub-task has a significant impact on the ability to produce high-quality summaries. Identifying important related sentences of text is a vital part of the extraction-based summarization method. The use of fuzzy logic as a sub-task for summarization greatly increased the consistency of the description. In the reference tables, the findings are easily noticeable. Compared to the performance of two online summarizers, our method generates stronger results. Thus, by integrating latent semantic analysis into the sentence function extracted fuzzy logic scheme to extract the semantic connections between concepts in the original document, our proposed approach increases the consistency of the summary. Although the emphasis of this paper is narrow: text summarization, the concepts are more widely relevant. In the future, we need to expand the proposed approach for multi-document summarization to accommodate broad data sets and domain-specific information.

REFERENCES

1. R. Nallapati, F. Zhai, and B. Zhou, "Summarunner: A recurrent neural network-based sequence model for extractive summarization of documents," in *Proceedings of the AAAI Conference on Artificial Intelligence*, 2017, vol. 31, no. 1.
2. C. A. Colmenares, M. Litvak, A. Mantrach, and F. Silvestri, "Heads: Headline generation as sequence prediction using an abstract feature-rich space," 2015.

3. R. Paulus, C. Xiong, and R. Socher, "A deep reinforced model for abstractive summarization," *arXiv preprint arXiv:1705.04304*, 2017.
4. P. Kouris, G. Alexandridis, and A. Stafylopatis, "Abstractive Text Summarization: Enhancing Sequence-to-Sequence Models Using Word Sense Disambiguation and Semantic Content Generalization," *Computational Linguistics*, vol. 47, no. 4, pp. 813-859, 2021.
5. L. Liu, Y. Lu, M. Yang, Q. Qu, J. Zhu, and H. Li, "Generative adversarial network for abstractive text summarization," in *Proceedings of the AAAI Conference on Artificial Intelligence*, 2018, vol. 32, no. 1.
6. M. Moradi and N. Ghadiri, "Different approaches for identifying important concepts in probabilistic biomedical text summarization," *Artificial intelligence in medicine*, vol. 84, pp. 101-116, 2018.
7. C. Mallick, A. K. Das, W. Ding, and J. Nayak, "Ensemble summarization of bio-medical articles integrating clustering and multi-objective evolutionary algorithms," *Applied Soft Computing*, vol. 106, p. 107347, 2021.
8. J. Grimmer, M. E. Roberts, and B. M. Stewart, "Machine learning for social science: An agnostic approach," *Annual Review of Political Science*, vol. 24, pp. 395-419, 2021.
9. [9] M. Mridha, A. A. Lima, K. Nur, S. C. Das, M. Hasan, and M. M. Kabir, "A Survey of Automatic Text Summarization: Progress, Process and Challenges," *IEEE Access*, vol. 9, pp. 156043-156070, 2021.
10. H. Balinsky, A. Balinsky, and S. Simske, "Document sentences as a small world," in *2011 IEEE International Conference on Systems, Man, and Cybernetics*, 2011, pp. 2583-2588: IEEE.
11. G. Erkan and D. R. Radev, "Lexrank: Graph-based lexical centrality as salience in text summarization," *Journal of artificial intelligence research*, vol. 22, pp. 457-479, 2004.
12. V. JAIN and S. PRASAD, "EVALUATION AND VALIDATION OF ONTOLOGY USING PROTEGE TOOL," *International Journal of Research in Engineering & Technology*, vol. 4, no. 4, pp. 21-32.
13. M. Hu, P. Liu, L. Bo, Y. Mao, K. Xu, and W. Su, "ProPC: A Dataset for In-Domain and Cross-Domain Proposition Classification Tasks," in *CCF International Conference on Natural Language Processing and Chinese Computing*, 2021, pp. 53-64: Springer.
14. A. Nenkova, L. Vanderwende, and K. McKeown, "A compositional context sensitive multi-document summarizer: exploring the factors that influence summarization," in *Proceedings of the 29th annual international ACM SIGIR conference on Research and development in information retrieval*, 2006, pp. 573-580.
15. X. Zhang, F. Wei, and M. Zhou, "HIBERT: Document level pre-training of hierarchical bidirectional transformers for document summarization," *arXiv preprint arXiv:1905.06566*, 2019.
16. J. Kupiec, J. Pedersen, and F. Chen, "A trainable document summarizer," in *Proceedings of the 18th annual international ACM SIGIR conference on Research and development in information retrieval*, 1995, pp. 68-73.
17. J. M. Conroy and D. P. O'leary, "Text summarization via hidden markov models," in *Proceedings of the 24th annual international ACM SIGIR conference on Research and development in information retrieval*, 2001, pp. 406-407.
18. R. Mihalcea, "Language independent extractive summarization," in *ACL*, 2005, vol. 5, pp. 49-52.
19. S. Li, Y. Ouyang, W. Wang, and B. Sun, "Multi-document summarization using support vector regression," in *Proceedings of DUC*, 2007: Citeseer.

20. M. Galley, "A skip-chain conditional random field for ranking meeting utterances by importance," 2006.
21. R. Barzilay and K. R. McKeown, "Sentence fusion for multidocument news summarization," *Computational Linguistics*, vol. 31, no. 3, pp. 297-328, 2005.
22. C.-S. Lee, Z.-W. Jian, and L.-K. Huang, "A fuzzy ontology and its application to news summarization," *IEEE Transactions on Systems, Man, and Cybernetics, Part B (Cybernetics)*, vol. 35, no. 5, pp. 859-880, 2005.
23. H. Tanaka, A. Kinoshita, T. Kobayakawa, T. Kumano, and N. Kato, "Syntax-driven sentence revision for broadcast news summarization," in *Proceedings of the 2009 Workshop on Language Generation and Summarisation (UCNLG+ Sum 2009)*, 2009, pp. 39-47.
24. P.-E. Genest and G. Lapalme, "Fully abstractive approach to guided summarization," in *Proceedings of the 50th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, 2012, pp. 354-358.
25. M. Dey and D. Das, "A Deep Dive into Supervised Extractive and Abstractive Summarization from Text," in *Data Visualization and Knowledge Engineering*: Springer, 2020, pp. 109-132.
26. P.-E. Genest and G. Lapalme, "Framework for abstractive summarization using text-to-text generation," in *Proceedings of the workshop on monolingual text-to-text generation*, 2011, pp. 64-73.
27. I. F. Moawad and M. Aref, "Semantic graph reduction approach for abstractive Text Summarization," in *2012 Seventh International Conference on Computer Engineering & Systems (ICCES)*, 2012, pp. 132-138: IEEE.