

A Novel Approach for Graph-based Extractive Text Summarization using Karıcı Dominant Set Algorithm and Eigenvector Centrality

Taner UÇKAN¹, Abdulsamet AYDIN²

¹Van Yüzüncü Yıl University, Department of Computer Engineering, 65000 Van, Turkey

²Van Yüzüncü Yıl University, Department of Artificial Intelligence and Robotics, 65000 Van, Turkey

✉: taneruckan@yyu.edu.tr  [0000-0001-5385-6775](https://orcid.org/0000-0001-5385-6775)  [0000-0002-5329-4407](https://orcid.org/0000-0002-5329-4407)

Received (Geliş): 10.04.2025

Revision (Düzeltilme): 14.05.2025.

Accepted (Kabul): 31.05.2025

ABSTRACT

This study presents a novel contribution to graph-based text summarization by integrating the Karıcı Dominant Clustering Algorithm into summarization systems for the first time. In the proposed method, a neighborhood matrix based on the number of shared words between sentences is used to construct the initial graph. The Karıcı Dominant Clustering Algorithm is then applied to identify dominant clusters, and the corresponding sentences are removed from the text. A second graph is constructed from the remaining sentences, and eigenvector centrality values are used to determine the most central sentences, which form the final summary. The method was evaluated on the DUC-2002 and DUC-2004 datasets using ROUGE metrics, achieving scores of 0.35748, 0.49049, and 0.57586 for 100-, 200-, and 400-word summaries, respectively. The experimental results demonstrate that the proposed approach outperforms several existing methods and provides a significant contribution to the field of automatic text summarization.

Keywords: Graph dominating set, Graph-based document summarization, Generic document summarization, Extractive text summarization, Multi document text summarization

Karıcı Baskın Küme Algoritması ve Özvektör Merkeziliği Kullanarak Çizge Tabanlı Çıkarımsal Metin Özetleme için Yeni Bir Yaklaşım

ÖZ

Bu çalışma, çizge tabanlı özetleme yöntemlerine yenilikçi bir katkı sunarak Karıcı Baskın Kümeleme Algoritması'nı ilk kez metin özetleme sistemlerine entegre etmektedir. Önerilen yöntemde, özetlenecek metindeki cümleler arasında ortak kelime sayılarına dayalı bir komşuluk matrisi oluşturularak ilk çizge inşa edilir. Ardından, Karıcı Baskın Kümeleme Algoritması kullanılarak çizgedeki baskın kümeler belirlenir ve bu kümeye ait cümleler metinden çıkarılır. Kalan cümlelerle ikinci bir çizge oluşturulur ve bu çizgedeki özvektör merkezilik değerlerine göre en merkezi cümleler seçilerek özet oluşturulur. Yöntem, DUC-2002 ve DUC-2004 veri kümeleri üzerinde ROUGE metrikleriyle değerlendirilmiş ve sırasıyla 100, 200 ve 400 kelimelik özetler için 0.35748, 0.49049 ve 0.57586 ROUGE değerlerine ulaşmıştır. Elde edilen bulgular, önerilen modelin mevcut yöntemlere kıyasla yüksek performans gösterdiğini ve literatüre anlamlı bir katkı sunduğunu ortaya koymaktadır.

Anahtar Kelimeler: Çizge baskın küme, Çizge tabanlı belge özetleme, Genel belge özetleme, Çıkarımsal metin özetleme, Çoklu belge metin özetleme

INTRODUCTION

With the rapid development of internet technology, it is seen that the information on the internet has increased at an extraordinary rate [1], [2]. The exponential growth of the information produced every day makes it difficult for people to access the right information from this information stack. During access to information, many documents related to the information sought can be obtained, but it is a laborious and time-consuming task to determine whether the requested information is actually found in the documents. Text summarization systems are being developed to help people save time, increase productivity, and access information more easily. Automatic document summarization is an area where the

issue has yet to be resolved despite recent advances. To increase the text's value for readers, various strategies are employed to extract its key information effectively [3]. For this reason, researchers are working on advanced or new methods to provide more efficiency regarding automatic text summarization technologies [4]. Summarization aims to transform extensive and detailed textual content into a more concise format, without compromising the core meaning and essential points. Text summarization aims to produce a concise and logically structured representation of the original text. A range of approaches has been developed including extractive, abstractive, and hybrid summarization methods, each offering distinct advantages and limitations [5]. Regardless of the approach employed,

text summarization is essential for enhancing the comprehension and distribution of information in an era characterized by rapid information flow and high data volume.

Extractive summarization involves constructing the summary by identifying and directly using the most relevant sentences from the original document. Abstractive text summarization involves generating summaries by producing new sentences that do not appear in the original text, following a comprehensive analysis of the document. [6]. In the proposed extractive text summarization method, dominant sets [7] and graph theory techniques are used.

In the subsequent sections of the study, Chapter 3 outlines the overall structure of the proposed text summarization approach, including the identification of dominant clusters within the graphs. Chapter 4 presents the experimental results, offering a detailed description in terms of the datasets and performance measures employed. Moreover, a systematic comparison is performed in comparison with the proposed model and leading state-of-the-art methods in the field. Finally, Chapter 5 offers a critical discussion and interpretation of the experimental results, highlighting their implications and relevance within the broader context of the study.

RELATED WORK

In text summarization research, numerous criteria have been incorporated into sentence scoring to enhance the quality of system-generated summaries. These include factors such as the sentence's position within the document, its length, the identification of key terms, the overlap of sentence content with the document's title, the presence of numerical values, and the frequency of terms. Integrating these features into the evaluation process has significantly improved summarization performance [8], [9], [10]. [11] proposes a graph-based summarization system using both BM25+ and TextRank algorithms. In this summarization system, a graph was obtained with the nodes of the sentences and the similarity score between the two sentences with the edge weights. Nodes with the maximum order are selected and summarized. [12] proposed a single-document graph-based method for generating a coherent summary from Arabic texts. The source text was first transformed into a textual graph and both statistical and semantic factors were used to evaluate each sentence according to a new method that takes into account relevance, scope and variety. A subgraph is then constructed to reduce the overall size of the document. Finally, the final summary was produced by excluding unnecessary and less significant expressions from the summarized sentences. [13]. Ranked for each sentence by calculating subject information, semantic content, important keywords and location features. The final scores of each sentence in the document were revealed by combining the ranking values calculated for each feature. The top-ranked sentences are identified and integrated into the final summary. [14] proposed a graph-based approach that consider the similarity between

sentences and each other and between sentences and the whole document. In this approach, the similarities of the sentences to each other and the similarities of the sentences to the subject of the document were evaluated together and weighted graphs were obtained. A summary is created by determining the nodes with the top-ranked values in the graph and including the sentences associated with them. Karci Summarization focuses on entropy centrality and a general summarization is made using various α values in Karci entropy calculation. Here, according to the calculated entropy result, the nodes that are considered to carry the most information are selected and summarized. [15] [16], introduced a graph-based summarization approach that leverages biomedical-specific knowledge along with a data mining technique known as frequent itemset mining. In their studies, they used the Jaccard Similarity approach to specify the similarities of the sentences. In the study on graph-based extractive text summarization, firstly, the effect of 4 different association methods on the "TextRank" method was investigated. DUC and CAST data sets were used as data sets. In addition to this study, a system has been developed using hierarchical associative clustering and "TextRank" methods. In the proposed method, sentences are clustered according to a certain criterion, and "TextRank" is utilized to extract relevant sentences from the clusters. According to the studies, it has been determined that the proposed system works better when DUC 2002 is used. In the CAST dataset, it has been determined that 2 methods out of 4 different association methods have passed, and the difference between the other 2 methods is small [17]. Regarding the graphs, dominant set, node coverage, perfect matching, max faction, node coloring etc. There are such problems. Obtaining the minimum dominant set is an NP-Hard problem [7]. Numerous studies have been carried out on the minimum dominant cluster. [18], tried to solve the minimum dominant cluster problem by using a two-stage ant colony (ACO) algorithm. [19], reported that summaries obtained using dominant clusters obtained more successful results than many methods. [20], applied the minimum dominant set problem to Hopfield networks and obtained successful results. In addition, they claimed that the algorithm they developed would also consider solutions to problems such as max clique and max independent set

PROPOSED SUMMARIZATION METHOD

The steps of the proposed approach for text summarization are illustrated by the block diagram in Figure 1. In this study, an extractive and multi-document summarization method using dominant sets in graphs is presented. In the study, texts with model summaries of 200 and 400 words from the DUC 2002 data set and texts with a 100-word model summary from the DUC 2004 data set were used as inputs to the proposed text summarization system. The proposed summarization technique is divided into three fundamental tasks. In the first task, static words (such as pronouns, prepositions,

conjunctions) and unnecessary characters are removed from the documents. It is aimed to increase the success rate by removing the words (stopwords) that do not make sense on their own. In the second stage, mathematical modeling and graph theory are used to visualize and quantify the semantic connections between sentences. Furthermore, this task incorporates identifying the nodes that make up the dominant cluster and extracting the sentences represented by these nodes from the original

text. A new graph is created with the remaining sentences in the text. In the last stage, the sentences corresponding to the nodes in the newly created graph are scored using the eigenvector centrality method. At this step of the suggested method, summaries of 100, 200 and 400 words were created from the texts, starting from the most important node. The effectiveness of the approach was subsequently evaluated in detail using various ROUGE performance metrics.

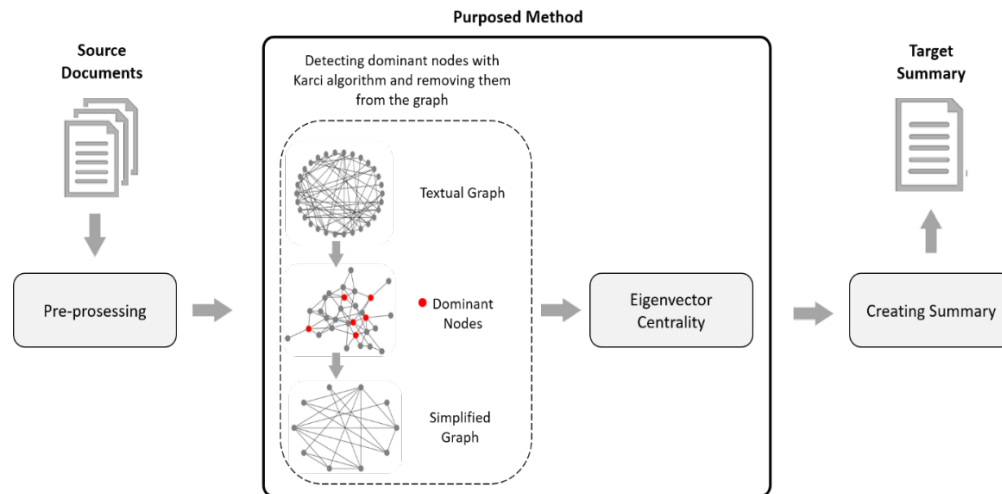


Figure 1. Schematic representation of the proposed document summarization model

According to graph theory, the dominant set is a subset of all the nodes in a graph. The vertices in the graph are either in the dominant cluster or are neighbors of at least one node in the dominant cluster [7]. Figure 2 illustrates a sample graph consisting of 10 nodes. As seen in Figure 3, dominant nodes dominate the entire graph.

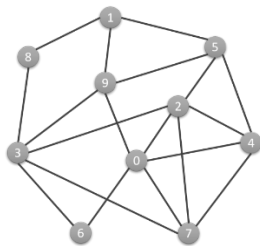


Figure 2. An example graph

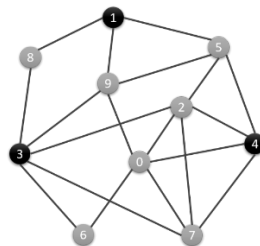


Figure 3. Dominant nodes of the graph

Detection of the dominant set

Within the scope of this study, the dominant set of nodes belonging to the textual graphs was determined using the Karci centrality value approach. This centrality algorithm is determined by calculating the node degrees, the Kmax tree degrees, and the basic cutoff degrees. This algorithm is composed of three fundamental stages. The initial step involves calculating the degree of each vertex in the graph, which represents the number of its connections to other nodes.

Kmax Tree

In the second stage, the degrees of the nodes in the tree were calculated using the Kmax Tree defined by Karci. The Kmax tree is a traversal of all the vertices in the graph, commencing from the vertices with the highest rank [21]. In the Kmax tree, the node possessing the highest degree is selected as the root node, and its adjacent nodes are enqueued according to their index sequence. When choosing the next highest ranked node among the neighboring nodes in the queue, if two or more nodes are the highest, the level of the nodes in the tree is checked first. Of these nodes, the node with the lowest level in the tree is given priority. If the equality continues, the first node is selected as the highest ranked node, taking into account its order in the queue. Kmax tree is obtained by adding all the nodes in the graph to the tree. Algorithm 1 shows the algorithm used to create the Kmax tree

Algorithm 1. Generating the Kmax Tree

```

1  Function Kmax(graph,root)
2      T = New Graph
3      visited, queue = empty list
4      add root to visited and queue list
5      while the queue is not empty
6          find root's neighbours

```

- 7 remove links between root and its neighbours
- 8 if the neighbour has not been visited before,
add it to the queue and visited list
- 9 if the neighbour has been visited before, it
comes out of the line
- 10 if there is a node with degree 0 in the queue,
exit the queue
- 11 root = highest degree node remaining in the
graph
- 12 return T

While constructing the Kmax tree of the graph in Figure 2 by using Algorithm 1, the highest order node 0 was chosen as the root. Then, nodes 2,4,6,7 and 9, which are neighbors of node 0, are added to the root as children. If there is a connection between these nodes, this connection is broken and the node degrees are updated. The reason for this is that since a node is reached over an edge, other alternative routes are no longer needed. According to the updated node degrees, node 9 was chosen as the highest degree node. By repeating the previous steps, nodes adjacent to node 9 are added to the tree. In the next step, node 1 became the highest order node. Neighbors to this node are also added to the tree. Thus, since all nodes are reached, the Kmax tree is created. Figure 4 shows the Kmax tree of the graph.

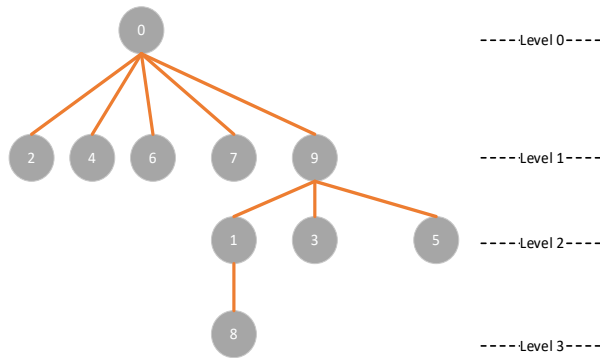
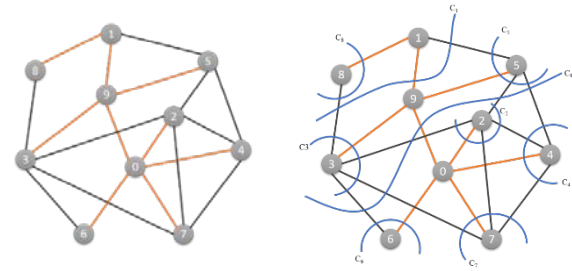


Figure 4. Kmax tree of the graph
Calculation of Fundamental Cutoffs

In the third step of the Karcı centrality algorithm, the basic cutting degrees are calculated. To calculate these basic cutting degrees, all the edges of the Kmax tree are cut out one by one and the graph is divided into two subsets. The connections of the nodes in the created clusters with the opposite cluster are determined and the cut-off value of which node has a connection is added to it. The basic cutting set is obtained as much as the number of branches in the Kmax tree [21]. Figure 5(a) shows the intersection of the sample graph and the Kmax tree. Figure 5(b), on the other hand, shows all the basic cuts performed according to the generated Kmax tree. In each cutting operation, only one edge specified in the Kmax tree is cut.



(a) Intersection of example graph and Kmax tree (b) All key cuts in the graph

Figure 5. Kmax tree of the example graph and all cuts

For example, $C_8 = \{\{8\}, \{0,1,2,3,4,5,6,7,9\}\}$ nodes are divided into two different clusters with the basic cut C_8 shown in Figure 5(b). When the connections between these two clusters are examined, it has been determined that there are edges between the nodes (8,1) and (8,3). In this cutting process, node 8 is given 2 cutting degrees, and nodes 1 and 3 are given 1 cutting degree each. In this way, all cuts are performed and the cutoff degree of each node is calculated. In Algorithm 2, the algorithm used to calculate the basic cutoff degrees is shown.

Algorithm 2. Calculation of Fundamental Cutoffs

```

1  Function CutSet(graph,tree)
2      treeNodes, treeEdges = sorted list of nodes and
3      edges in tree)
4      set1, set2 = empty list
5      cut_set = empty dictionary
6      For each i in treeNodes
7          cut_set[i] = 0
8      X = array of zeros (length of treeEdges, length of
9      treeNodes)
10     If the size of X > 0
11         cut_set = dict(zip(treeNodes,[int(n) for n in
12         list(X[0])]))
13         For i ranging from 0 to (length of treeEdges-
14         1)
15             clear set1 and set2
16             T2 = copy of tree
17             Remove edges from T2 with [treeEdges[i]]
18             d=
19             sorted(list(nx.connected_components(T2)))
20             Expand set1 with d [0]
21             Expand set2 with d [1]
22             For each a in set1
23                 For each b in set2
24                     If graph has an edge between a and
25                     b
26                         Increment cut_set[a] and
27                         cut_set[b] by 1
28     Return cut_set

```

Karcı dominant set algorithm

Karcı- centrality algorithm is expressed with the node dominance value (Γ) symbol and the dominance value for each vertex in the graph is calculated by equation 1 [22].

$$\begin{aligned} \text{Node Dominance Value } (\Gamma) \\ = \text{Graph Degree} \\ + \text{Kmax Degree} \\ + \text{Cut Degree} \end{aligned} \quad (1)$$

After computing the dominance scores for each node in the graph, the subsequent step is to identify the dominant nodes. In Algorithm 3, the algorithm used to create the dominant cluster, and in Figure 6, the flow diagram of the Karcı dominant cluster algorithm is given. The first thing to be done during the determination of the dominant cluster is to determine whether there is a pendant node in the graph. Nodes with a degree of 1 are called pendant nodes [23]. If there is a pendant node in the graph, the node adjacent to this node is included in the dominant cluster. The node and its neighbors included in the dominant cluster are removed from the graph and the Kmax tree, and the graph is updated.

During the second phase, if no pendant nodes exist in the current graph, the node exhibiting the highest dominance value is added to the dominant cluster. If the dominance value of more than one node is equal, the node with the lower Kmax degree is selected. If the equality still exists, any of the nodes is added to the dominant cluster. In the next step, the dominant node and its neighbors are deleted from the graph and the Kmax tree and the graph is updated. In each iteration, pendant node control is performed. The dominant cluster is formed by iteratively performing the operations until all nodes are removed from both the graph and the Kmax tree.

Algorithm 3. Determination of Dominant Cluster

```

1  Function Dominating_Set(graph)
2    G = graph
3    T = Kmax tree
4    graphDegree = a dictionary of node and its degree
    in graph
5    kmaxDegree = a dictionary of node and its degree
    in T
6    cutDegree = a dictionary of node and its cut degree
    in graph
7    nodeDominanceValue = a dictionary of node and
    its dominance value in graph
8    pendantNodes = list to store pendant nodes
9    dominatingSet = list to store dominating nodes
10   while graph still has nodes
11     clear pendantNodes
12     for each node in graphDegree
13       if the degree of node is 1

```

```

14         add node to
            pendantNodes
15     if there are pendantNodes
16         select the first node in
            pendantNodes as the chosen node
17         find its neighbors in graph
18         select the first neighbor as the
            dominant node
19     else
20         find the vertex with the highest
            dominance value in graph
21         if there are multiple nodes with the
            highest value
22             find the node with the
                lowest kmax degree
                among them
23             select the first node with
                the lowest degree as the
                dominant node
24         else
25             select the first node with
                the highest dominance
                value as the dominant
                node
26         add dominantNode to dominatingSet
27         find the neighbors of dominantNode in graph
28         remove dominantNode and its neighbors
            from graph and T
29         update graphDegree, kmaxDegree,
            cutDegree and nodeDominanceValue
30     return dominatingSet

```

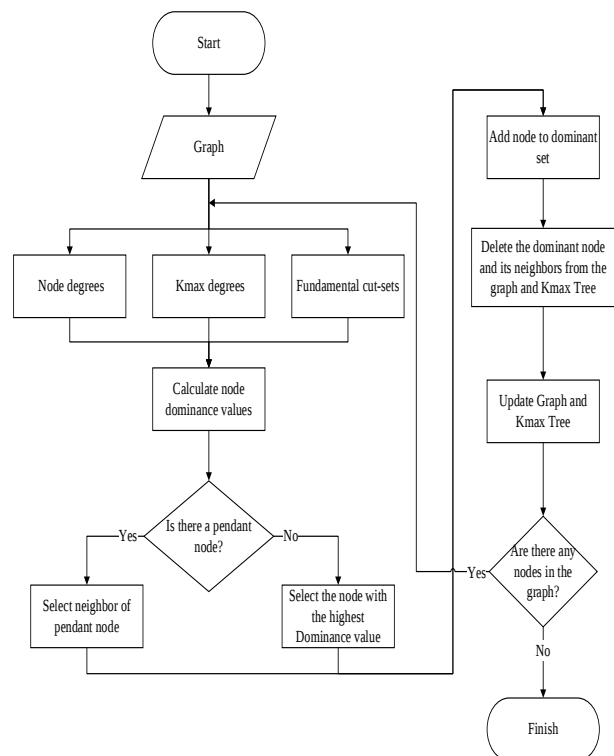


Figure 6. Flowchart of Karcı Dominant Set Algorithm

Eigenvector centrality

Eigenvector centrality is a technique employed to assess the significance of a node within a graph, relying on the notion that links to highly influential node contribute more to a node's importance than links to less influential ones [24]. The eigenvector centrality is calculated by equation 2, where V_j and V_i are the weights of the connections between the nodes and λ is a constant [25].

$$C_e(v_i) = \frac{1}{\lambda} \sum_{v_j \in N(v_i)} v_j C_e(v_j) \quad (2)$$

Demonstration of the Karci dominant set algorithm on DUC 2002

Due to the large size of the graphs obtained from the texts in the dataset, the graph obtained from the first ten sentences of the d061 file is shown in Figure 7 and the dominant nodes of this graph are shown in Figure 8. Table 1 shows the sentences corresponding to each node in the graph in Figure 7.

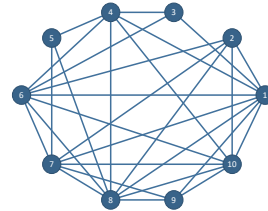


Figure 7. The graph obtained from the first 10 sentences in d061.txt

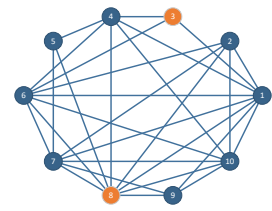


Figure 8. Dominant nodes of the graph obtained from the first 10 sentences in d061.txt

Table 1. First 10 sentences of file d061 of DUC 2002 [26]

1	Hurricane Gilbert swept toward the Dominican Republic Sunday, and the Civil Defense alerted its heavily populated south coast to prepare for high winds, heavy rains and high seas.
2	The storm was approaching from the southeast with sustained winds of 75 mph gusting to 92 mph.
3	"There is no need for alarm," Civil Defense Director Eugenio Cabral said in a television alert shortly before midnight Saturday.
4	Cabral said residents of the province of Barahona should closely follow Gilbert's movement.
5	An estimated 100,000 people live in the province, including 70,000 in the city of Barahona, about 125 miles west of Santo Domingo.
6	Tropical Storm Gilbert formed in the eastern Caribbean and strengthened into a hurricane Saturday night.
7	The National Hurricane Center in Miami reported its position at 2 am Sunday at latitude 161 north, longitude 675 west, about 140 miles south of Ponce, Puerto Rico, and 200 miles southeast of Santo Domingo.
8	The National Weather Service in San Juan, Puerto Rico, said Gilbert was moving westward at 15 mph with a "broad area of cloudiness and heavy weather" rotating around the center of the storm.
9	The weather service issued a flash flood watch for Puerto Rico and the Virgin Islands until at least 6 pm Sunday.
10	Strong winds associated with the Gilbert brought coastal flooding, strong southeast winds and up to 12 feet feet to Puerto Rico's south coast.

The DUC 2002 and DUC 2004 datasets include both the source texts and corresponding human-generated reference summaries. The recommended approach is assessed in comparison with other competitive methods using the reference summaries as a benchmark.

Extractive text summarization methods are not yet as natural as human-generated summaries since the sentences are taken as they are from the texts. Table 2 shows the model summary generated by humans and the 200-word summary obtained by the proposed system.

Table 2. 200-word model and system summaries

Model Summary ([26])	System Summary
Tropical Storm Gilbert formed in the eastern Caribbean and strengthened into a hurricane Saturday night. Hurricane Gilbert slammed into Kingston on Monday with torrential rains and 115 mph winds that ripped roofs off homes and buildings, uprooted trees and downed power lines. The storm killed 19 people in Jamaica and five in the Dominican Republic before moving west to Mexico. The Jamaican Embassy	Heavy rain and stiff winds downed power lines and caused flooding in the Dominican Republic on Sunday night as the hurricane's center passed just south of the Barahona peninsula, then less than 100 miles from neighboring Haiti. " Hurricane Gilbert, packing 110 mph winds and torrential rain, moved over this capital city today after skirting Puerto Rico, Haiti and the Dominican Republic. Hurricane Gilbert

reported earlier that 500,000 of the nation's 2.3 million people were homeless. Gilbert also buffeted the Cayman Islands, but no deaths were reported. Hurricane Gilbert, one of the strongest storms ever, slammed into the Yucatan Peninsula Wednesday and leveled thatched homes, tore off roofs, uprooted trees and cut off the Caribbean resorts of Cancun and Cozumel. The storm was about 550 miles southeast of Brownsville, Texas, the center said in a statement. The National Hurricane Center said Gilbert was the most intense storm on record in terms of barometric pressure. Earlier Wednesday Gilbert was classified as a Category 5 storm, the strongest and deadliest type of hurricane. Hurricane Gilbert's growth from a harmless low-pressure zone off Africa to a ferocious killer in the Gulf of Mexico was fueled by a combination of heat, moisture and wind that baffles forecasters.

swept toward the Dominican Republic Sunday, and the Civil Defense alerted its heavily populated south coast to prepare for high winds, heavy rains and high seas. The National Hurricane Center in Miami reported its position at 2 am Sunday at latitude 161 north, longitude 675 west, about 140 miles south of Ponce, Puerto Rico, and 200 miles southeast of Santo Domingo. Forecasters said the hurricane was gaining strength as it passed over the ocean and would dump heavy rain on the Dominican Republic and Haiti as it moved south of Hispaniola, the Caribbean Island they share, and headed west. " At midnight EDT Gilbert was centered near latitude 215 north, longitude 902 west and approaching the north coast of Yucatan, about 60 miles east-northeast of the provincial capital, Merida, the National Hurricane Center in Coral Gables, Fla, said. The storm ripped the roofs off houses and flooded coastal areas of southwestern Puerto Rico after reaching hurricane strength off the island's southeast Saturday night.

EXPERIMENTAL RESULTS

This study evaluates the proposed method's performance by employing the Document Understanding Conference (DUC) datasets, known for being among the most prevalent benchmark resources in text summarization studies [26]. DUC datasets are used to develop and test text summarization algorithms. DUC dataset is a dataset prepared for use in the text summarization field. The dataset contains a summary and a full text for each news article. With the proposed approach, abstracts of 200 and 400 words were obtained from the texts in the DUC 2002 dataset and 100-word summaries from the texts in the DUC 2004 dataset were obtained and compared with the model summaries. The DUC 2002 dataset comprises 59 document clusters, each consisting of approximately 10 news articles, totaling 567 documents. It has been observed that each document contains an average of 25 to 30 sentences. The DUC 2004 dataset includes 500 documents, each accompanied by human-generated reference summaries. The BBC News dataset is a dataset used for data categorization consisting of 2225 documents from 2004-2005, corresponding to stories in five domains. This dataset is designed for extractable text summarization purposes [27]. The "News Articles" folder contains 417 political news articles published by the BBC between 2004 and 2005. For each article, there are five different summaries in the "Summaries" folder. The BBC News dataset consists of 2,225 documents spanning five distinct categories: business, entertainment, politics, sport, and technology. Each article contains an average of 19 sentences, and multiple human-written reference summaries are available for each document.

To evaluate the effectiveness of the abstracts produced in this study, ROUGE (Recall-Oriented Understudy for Gisting Evaluation) performance metrics, which are among the most commonly applied evaluation criteria in

text summarization systems, were utilized. ROUGE counts the common n-grams (n-grams refers to n word groups in a text) of the summary text produced by text summarization systems and the abstract text created by humans, and the ratios of these common n-grams to the total n-grams [9]. The high score between 0 and 1 produced by the ROUGE criteria represents the success of the automatic summary. In this study, the performance of the suggested method is assessed using ROUGE-N, ROUGE-L, ROUGE-W-1.2 and ROUGE-SU performance measures.

$$ROUGE - N = \frac{\sum_{C \in \{ReferenceSummaries\}} \sum_{gram_n \in C} count_{match}(gram_n)}{\sum_{C \in \{ReferenceSummaries\}} \sum_{gram_n \in C} count(gram_n)}$$

In this study, the effectiveness of the generated summaries was evaluated using ROUGE (Recall-Oriented Understudy for Gisting Evaluation), a set of performance metrics commonly employed in text summarization systems. ROUGE-N measures the n-gram overlap between the system-generated summaries and the human-written reference summaries; ROUGE-L assesses the longest common subsequence; and ROUGE-SU evaluates similarity based on skip-bigrams and unigrams. An approach that has not been used in any summary study before was used in our study. Assuming that the sentences associated with the dominant nodes should be excluded from the abstract, a new graph was constructed using the remaining sentences, with those corresponding to the dominant nodes removed from the original text. The nodes of this new graph were scored using the eigenvector centrality method and summaries were obtained. In Table 3, 100-word summaries, in Table 4, 200-word summaries and in Table 5, 400-word summaries are reported separately according to the DUC

dataset, and in Table 6, summaries obtained from the BBC New dataset are reported separately according to Recall, Precision and F-Score using ROUGE performance metrics. The evaluation metrics employed in this study include Rouge-1, Rouge-2, Rouge-3, Rouge-4, Rouge-L, Rouge-

W-1.2, Rouge-S*, and Rouge-SU*. In Tables 3-6, the first column displays the various types of summarization performance metrics, while the subsequent columns present their average performance across Recall, Precision, and F-Score.

Table 3. Performance values of 100-word summaries (DUC 2004)

	Recall	Precision	F-score
Rouge-1	0.38691	0.33342	0.35748
Rouge-2	0.07818	0.06743	0.07227
Rouge-3	0.02378	0.02049	0.02197
Rouge-4	0.00677	0.00582	0.00625
Rouge-L	0.30356	0.26250	0.28099
Rouge-W-1.2	0.10327	0.16345	0.12630
Rouge-S*	0.12957	0.09561	0.10914
Rouge-SU*	0.13437	0.09944	0.11338

Table 4. Performance values of 200-word summaries (DUC 2002)

	Recall	Precision	F-score
Rouge-1	0.51702	0.46710	0.49049
Rouge-2	0.24703	0.22194	0.23369
Rouge-3	0.18815	0.16832	0.17759
Rouge-4	0.17194	0.15357	0.16215
Rouge-L	0.48475	0.43786	0.45984
Rouge-W-1.2	0.17469	0.24875	0.20512
Rouge-S*	0.23343	0.19014	0.20908
Rouge-SU*	0.23613	0.19253	0.21162

Table 5. Performance values of 400-word summaries (DUC 2002)

	Recall	Precision	F-score
Rouge-1	0.58847	0.56406	0.57586
Rouge-2	0.32113	0.30750	0.31409
Rouge-3	0.25830	0.24701	0.25247
Rouge-4	0.23775	0.22717	0.23228
Rouge-L	0.55874	0.53538	0.54666
Rouge-W-1.2	0.17626	0.26609	0.21197
Rouge-S*	0.31889	0.29281	0.30497
Rouge-SU*	0.32019	0.29407	0.30626

Table 6. Performance values of summaries (BBC News)

	Recall	Precision	F-score
Rouge-1	0.77617	0.63884	0.69058
Rouge-2	0.66749	0.55359	0.59606
Rouge-3	0.63031	0.52324	0.56306
Rouge-4	0.60492	0.50180	0.54007
Rouge-L	0.57108	0.47136	0.50919
Rouge-W-1.2	0.18591	0.40615	0.25057
Rouge-S*	0.46384	0.31341	0.35355
Rouge-SU*	0.46868	0.31740	0.35797

To determine the relative impact and importance of each component, an ablation analysis was carried out as part of our study. In this study, the Node Degree is first removed from the model. The performance of the resulting model is analysed using ROUGE metrics. Then,

the performance of the model was analysed by removing the Kmax Degree and the Cutting Degree from the model separately, respectively. As shown in Table 7, the ablation study reveals that every component tested is vital to the success of the proposed model, and omitting any

of them results in reduced performance. In this study, the DUC 2002 dataset and 200-word summaries were used.

Table 7. Results of the ablation study

Component	Rouge-1	Rouge-2	Rouge- L	Rouge-W1.2	Rouge- SU*
Node+Kmax Degree	0.47639	0.21202	0.44299	0.19315	0.20699
Node+Cut Degree	0.44858	0.18366	0.41557	0.17715	0.18603
Kmax+Cut Degree	0.46181	0.19409	0.43013	0.18704	0.19088
Dominance value	0.49049	0.23369	0.45984	0.20512	0.21162

DISCUSSION AND CONCLUSION

A comparative analysis was conducted between the abstracts produced by the proposed approach and those presented in earlier research.

Luhn [28] generated summaries by assigning scores to sentences based on statistical information derived from word frequencies and distributions, in combination with machine learning techniques. Landauer et al. [29], [30] introduced a novel approach that relies solely on general mathematical methods. Mihalcea utilized an extractive and unsupervised method known as TextRank [31], [32], where the text's linked structure was converted into graphs, and summaries were generated by scoring the sentences based on their significance. One of the key advantages of the TextRank algorithm is that it builds upon Google's PageRank [33], eliminating the necessity for a manually defined structure. In related work, Erkan et al. suggested a graph-based technique to evaluate sentence significance. In their approach, known as LexRank, sentence importance was determined by calculating the eigenvector centrality of the vertices representing the sentences in the graph [34]. SumBasic and KL-Sum are statistical methods used in summarizing text. In SumBasic, each sentence in the text is evaluated based on its importance, and the most significant ones are selected to build the summary. This importance score is

based on the total frequencies in the text of the words in each sentence. That is, the more frequently a word is used, the more important the sentence it is in is considered [35]. In the KL-Sum method, the resemblance of each sentence to other sentences in the text is calculated using the Kullback-Leibler (KL) distance measure. The weight of each sentence is determined by both the importance of its words and its degree of similarity to other sentences. The sentences with the highest weights are chosen to form the summary [36]. These methods were selected due to their reliance on mathematical, statistical, or graph-based approaches, which align closely with the methodology proposed in this study. Moreover, similar to the suggested approach, all of the chosen methods fall under the category of unsupervised document summarization.

The ROUGE performance scores of the Luhn, LSA, TextRank, LexRank, SumBasic, and KL-Sum techniques, as well as the summarization method introduced in this research, are displayed in Table 8 using the DUC-2002 dataset. The highest values in each row of the table are emphasized in bold. It is clearly seen that the proposed approach outperforms all competing methods on 200-word summaries. Additionally, Figure 9 presents a graphical comparison of the F-score values for the 200-word summaries generated by the proposed method and other competing approaches.

Table 8. Comparison of the proposed method with other similar methods (DUC 2002, 2004)

	Rouge- 1	Rouge- 2	Rouge- L	Rouge-W1.2	Rouge-SU
100 Words	Luhn	0.26733	0.03182	0.20354	0.06534
	LSA	0.29759	0.03983	0.24845	0.08820
	TextRank	0.36292	0.07338	0.27991	0.11328
	LexRank	0.31255	0.05292	0.25053	0.08884
	SumBasic	0.32808	0.05413	0.23680	0.09437
	KL-Sum	0.32924	0.06946	0.26789	0.10117
	Ours	0.35748	0.07227	0.28099	0.11338
200 Words	Luhn	0.45924	0.20160	0.43022	0.18556
	LSA	0.37342	0.08921	0.33162	0.12527
	TextRank	0.45868	0.16000	0.44289	0.19200
	LexRank	0.46997	0.18010	0.42787	0.19309
	SumBasic	0.45597	0.15757	0.39194	0.17548
	KL-Sum	0.37104	0.12018	0.33741	0.13098

	Ours	0.49049	0.23369	0.45984	0.20512	0.21162
400 Words	Luhn	0.57746	0.32654	0.54940	0.21675	0.29118
	LSA	0.46804	0.16141	0.42685	0.15464	0.20140
	TextRank	0.55179	0.24447	0.53926	0.20648	0.28307
	LexRank	0.54188	0.23382	0.50793	0.19550	0.26541
	SumBasic	0.51462	0.18542	0.45995	0.17378	0.23197
	KL-Sum	0.43490	0.16132	0.39924	0.15442	0.17688
	Ours	0.57586	0.31409	0.54666	0.21197	0.30626

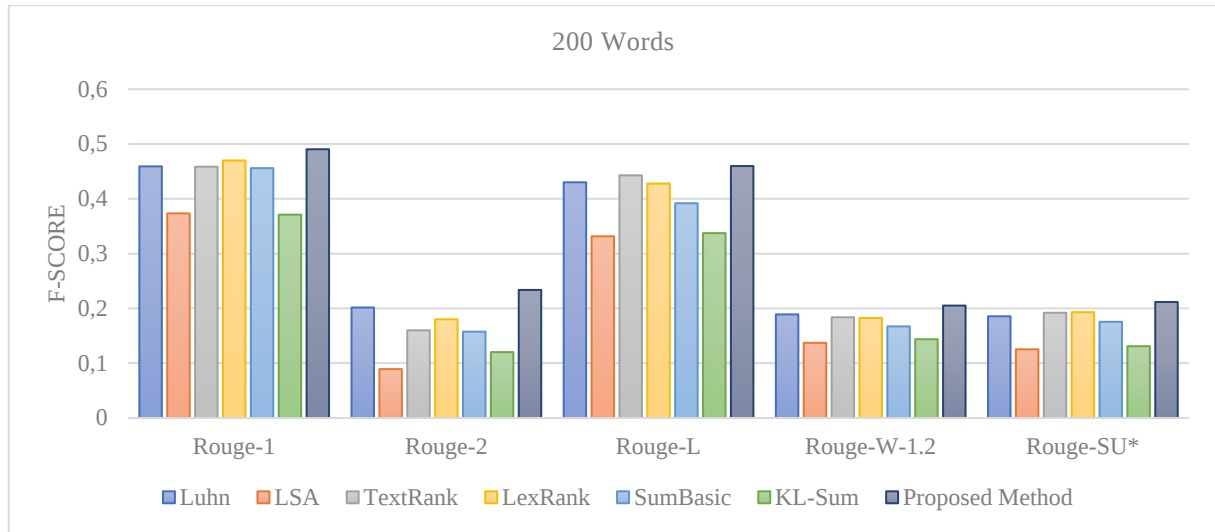


Figure 9. Graphical comparison of the proposed method with other competitive methods (200 words)

F-score values of 400-word summaries of all competitive methods, including the method suggested in Figure 10, were compared graphically.

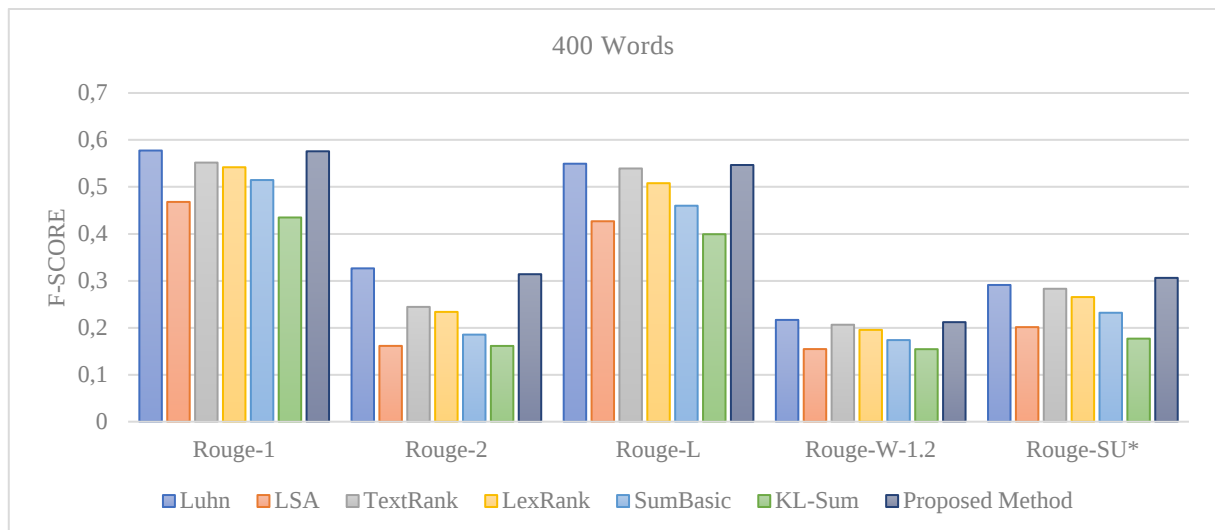


Figure 10. Graphical comparison of the proposed method with other competitive methods (400 words)

The comparison of F-score values across all methods for the 100-word summaries from the DUC-2004 dataset is displayed in Figure 11.

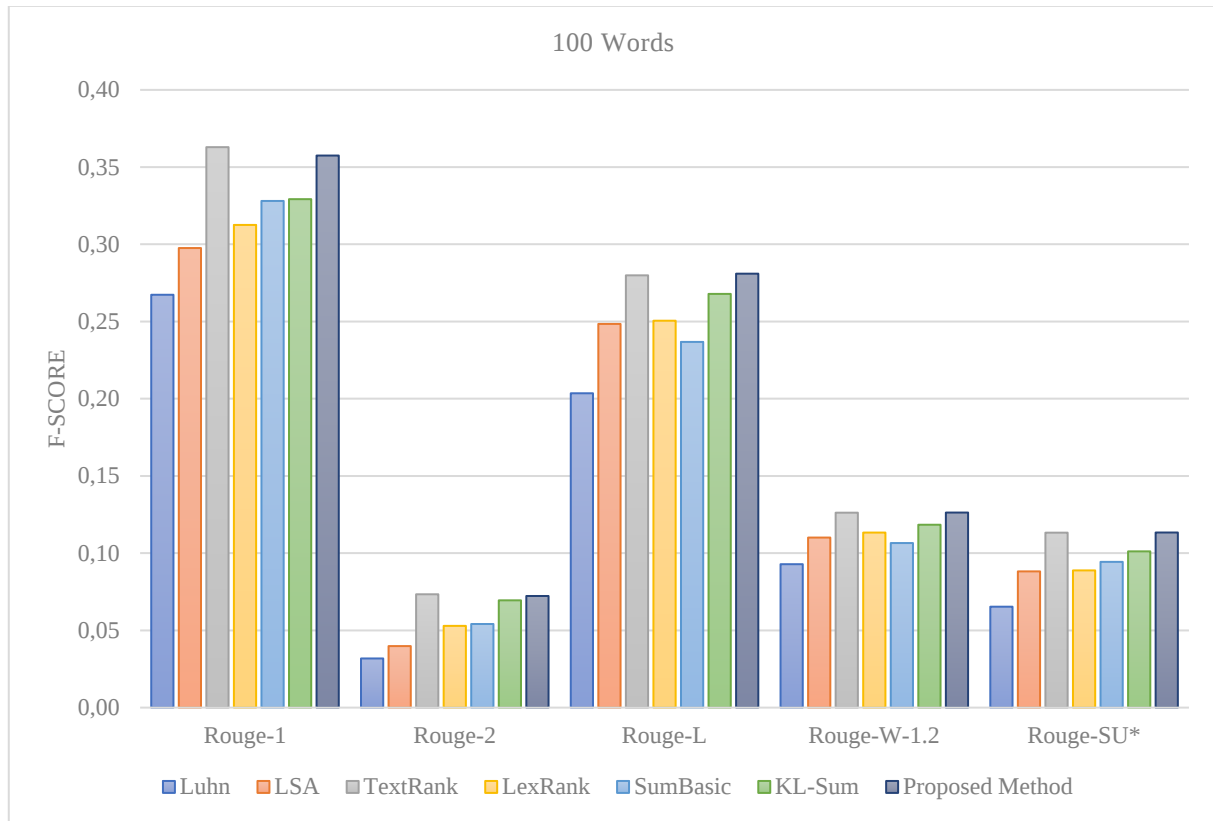


Figure 11. Graphical comparison of the suggested method with other competitive methods (100 words)

It is clearly seen that the proposed approach outperforms all competitive methods for 200-word summaries.

In 400-word summaries, it outperforms the Luhn method by 0.277% in Rouge-1 metric, 3.813% in Rouge-2 metric, 0.499% in Rouge-L metric, 2.205% in Rouge-W-1.2 metric, but 5.179% in Rouge-SU metric. Therefore, the suggested approach is superior to Luhn and other approaches.

For 100-word summaries, it is clearly seen that it outperforms all competitive approaches according to the F-score values of Rouge-L, Rouge-W-1.2 and Rouge-SU* metrics. When other metrics are also taken into

consideration, the TextRank method has shown superior performance compared to other approaches.

Table 9 illustrates that the proposed method outperforms all competitive approaches according to Rouge-1, Rouge-2, Rouge-L, and Rouge-W1.2 metrics. The proposed method outperforms the closest approach by 0.818% with respect to Rouge-1 metric, 2.497% with respect to Rouge-2 metric, 18.596% with respect to Rouge-L metric and 20.629% with respect to Rouge-W-1.2 metric. It underperformed the Rouge-SU metric by 7.233%. The graph comparing the F-score values of all methods for the summaries obtained from the BBC News dataset is shown in Figure 12.

Table 9. Comparison of the suggested method with other similar methods (BBC News)

		Rouge- 1	Rouge- 2	Rouge-L	Rouge-W1.2	Rouge-SU
BBC News	Luhn	0.68498	0.58154	0.42935	0.20772	0.38588
	LSA	0.54435	0.40154	0.37586	0.1713	0.24574
	TextRank	0.63133	0.51513	0.42793	0.20507	0.35254
	LexRank	0.67756	0.55625	0.48311	0.23147	0.38742
	SumBasic	0.56753	0.40593	0.37357	0.17317	0.26762
	KL-Sum	0.44298	0.28447	0.30028	0.12984	0.16749
	Ours	0.69058	0.59606	0.50919	0.25057	0.35797

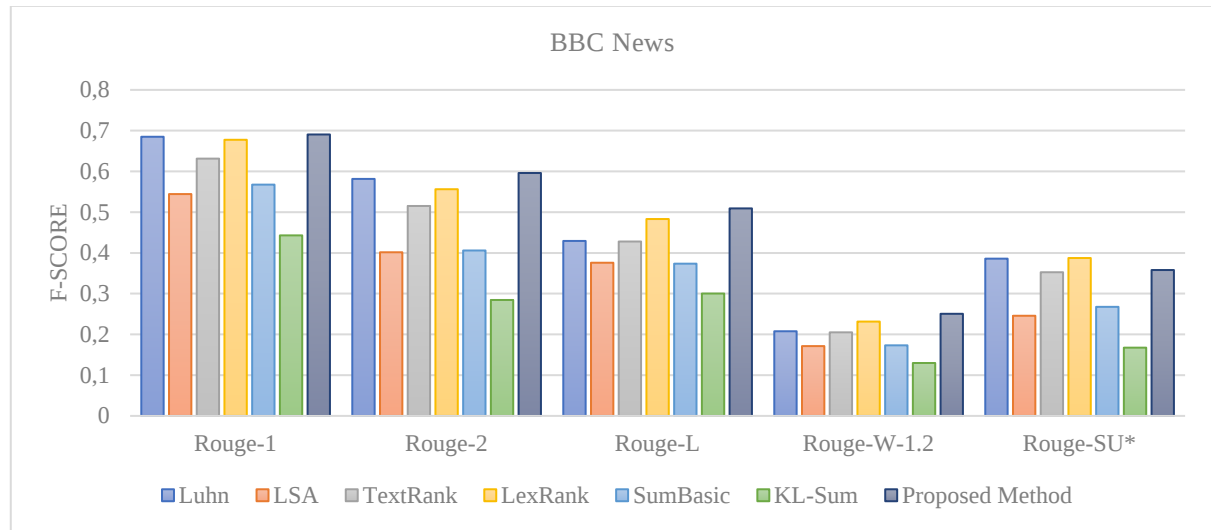


Figure 12. Graphical comparison of the proposed method with other competitive methods (BBC News)

Numerous investigations have explored various aspects of extractive approaches to text summarization [37]. Some of them are NN-SE [38], SummaRuNNer [39], Egraph [40], Tgraph [41], FbTS [42], FPGAC [43], GDSCO [44], SRL-ESA [45].

The values shown in Table 10 show that the suggested method better performance than all other methods on the

DUC 2002 dataset as measured by the ROUGE-1 and ROUGE-L metrics. Artificial neural networks generally need training data. They are also slow in the training and implementation phase [5]. The suggested method is an unsupervised summarization approach that is both language- and domain-independent.

Table 10. Comparison of the proposed model with other approaches (DUC 2002)

Model	Method	Rouge- 1	Rouge- 2	Rouge- L	Rouge-SU
FbTS	Optimization-Based	0.4782	0.2295	0.3362	0.2122
FPGAC	Fuzzy-Logic-Based	0.4868	0.2291	-	-
GDSCO	Optimization Based	0.4901	0.2304	-	-
SRL-ESA	Graph-Based	0.4620	0.2160	0.3070	0.2320
NN-SE	Neural Networks Based	0.4740	0.2300	-	-
SummaRuNNer	Neural Networks Based	0.4704	0.2400	0.1470	-
Egraph+coh	Graph-Based	0.4709	0.2380	-	-
Tgraph	Graph-Based	0.4810	0.2430	-	-
Proposed Method	Graph-Based	0.4904	0.2336	0.4598	0.2116

A comprehensive assessment of the proposed text summarization method has been carried out using various datasets and evaluation criteria. The results emphasize that the innovative document summarization approach, which is grounded in the dominant cluster of the graphs, produces superior outcomes. This is evident as it outperformed all the compared methods in the 200-word summaries and surpassed numerous methods in both the 400-word and 100-word summaries.

It is very important to save time in order to reach the right information from the enormously increasing data today. In order to find a solution to this problem, studies on automatic text summarization are increasing day by day and new methods are being developed. In this study, unlike graph-based text summarization methods based on sentence scoring, a graph-based text summarization system was developed by using dominant clusters.

In the proposed approach, the algorithm used to detect dominant clusters was proposed by Karcı [7]. By using this algorithm, a new graph was created by selecting the most dominant nodes in the graph and removing the sentences it represents from the document to be summarized. The eigenvector centrality of the vertices of the new graph is calculated and the sentences represented are added to the summary in order of significance.

To demonstrate the model's performance, Recall, Precision, and F-Scores were evaluated using ROUGE performance measures. The experiments were conducted multiple times with abstracts of 100, 200, and 400 words. The success of the proposed system has been verified with the results obtained during the experimental procedures using the DUC-2002, DUC-2004 and BBC News datasets.

The proposed summarization method achieved superior success for 200-word abstracts compared to all the methods compared. At the same time, it showed a high performance in 100 and 400-word summaries, leaving most of them behind.

Regarding performance, the suggested summarization method is anticipated to yield promising outcomes and motivate future research as a strong contender in the field of text summarization.

REFERENCES

- [1] M. R. Amini, N. Usunier, and P. Gallinari, 'Automatic text summarization based on word-clusters and ranking algorithms', *Lecture Notes in Computer Science*, vol. 3408, pp. 142–156, 2005, doi: 10.1007/978-3-540-31865-1_11/COVER.
- [2] A. Khan, N. Salim, and Y. Jaya Kumar, 'A framework for multi-document abstractive summarization based on semantic role labelling', *Appl Soft Comput*, vol. 30, pp. 737–747, May 2015, doi: 10.1016/J.ASOC.2015.01.070.
- [3] L. Ermakova, J. V. Cossu, and J. Mothe, 'A survey on evaluation of summarization methods', *Inf Process Manag*, vol. 56, no. 5, pp. 1794–1814, 2019.
- [4] H. Cengiz, T. Uçkan, E. Seyyarer, and A. Karcı, 'Graph-based suggestion for text summarization', in *2018 International Conference on Artificial Intelligence and Data Processing (IDAP)*, Ieee, 2018, pp. 1–6.
- [5] W. S. El-Kassas, C. R. Salama, A. A. Rafea, and H. K. Mohamed, 'Automatic text summarization: A comprehensive survey', 2021. doi: 10.1016/j.eswa.2020.113679.
- [6] T. Uçkan and C. Hark, 'Çıkarımsal metin özetleme yöntemleri', in *Endüstride Dijitalleşme Örnekleri*, S. Güldal, Ed., Ankara/ Turkey: Iksad Publications, 2022, pp. 31–46.
- [7] A. Karcı, 'New algorithms for minimum dominating set in any graphs', *Computer Science*, vol. 5, no. 2, pp. 62–70, 2020.
- [8] H. P. Edmundson, 'New methods in automatic extracting', *Journal of the ACM (JACM)*, vol. 16, no. 2, pp. 264–285, 1969.
- [9] C. Y. Lin, 'Rouge: A package for automatic evaluation of summaries', *Proceedings of the workshop on text summarization branches out (WAS 2004)*, no. 1, 2004.
- [10] G. Salton and C. Buckley, 'Term-weighting approaches in automatic text retrieval', *Inf Process Manag*, vol. 24, no. 5, pp. 513–523, 1988.
- [11] V. Gulati, D. Kumar, D. E. Popescu, and J. D. Hemanth, 'Extractive article summarization using integrated TextRank and BM25+ algorithm', *Electronics (Basel)*, vol. 12, no. 2, p. 372, 2023.
- [12] Y. A. AL-Khassawneh and E. S. Hanandeh, 'Extractive Arabic text summarization-graph-based approach', *Electronics (Basel)*, vol. 12, no. 2, 2023, doi: 10.3390/electronics12020437.
- [13] A. Joshi, E. Fidalgo, E. Alegre, and R. Alaiz-Rodriguez, 'RankSum—An unsupervised extractive text summarization based on rank fusion', *Expert Syst Appl*, vol. 200, p. 116846, 2022.
- [14] R. C. Belwal, S. Rai, and A. Gupta, 'A new graph-based extractive text summarization using keywords or topic modeling', *J Ambient Intell Humaniz Comput*, vol. 12, no. 10, pp. 8975–8990, 2021.
- [15] C. Hark and A. Karcı, 'Karcı summarization: A simple and effective approach for automatic text summarization using Karcı entropy', *Inf Process Manag*, vol. 57, no. 3, p. 102187, 2020.
- [16] M. N. Azadani, N. Ghadiri, and E. Davoodijam, 'Graph-based biomedical text summarization: An itemset mining and sentence clustering approach', *J Biomed Inform*, vol. 84, pp. 42–58, 2018.
- [17] C. Yalkın, 'Çizge tabanlı metin özetleme', Yüksek Lisans Tezi, Yıldız Teknik Üniversitesi Fen Bilimleri Enstitüsü, İstanbul, 2014.
- [18] R. Jovanovic and M. Tuba, 'Ant colony optimization algorithm with pheromone correction strategy for the minimum connected dominating set problem', *Computer Science and Information Systems*, vol. 10, no. 1, pp. 133–149, 2013.
- [19] C. Shen and T. Li, 'Multi-document summarization via the minimum dominating set', in *Proceedings of the 23rd International Conference on Computational Linguistics (Coling 2010)*, 2010, pp. 984–992.
- [20] X. Xu et al., 'An algorithm for the minimum dominating set problem based on a new energy function', in *SICE 2004 Annual Conference*, IEEE, 2004, pp. 924–926.
- [21] F. Öztemiz and A. Karcı, 'A New Approach to Determining Effective Nodes in Linked Graphs', *Dokuz Eylül Üniversitesi Mühendislik Fakültesi Fen ve Mühendislik Dergisi*, vol. 24, no. 70, pp. 143–155, 2021, doi: 10.21205/deufmd.2022247014.
- [22] F. Öztemiz, 'Karmaşık ağlarda hakim düğümlerin belirlenmesi için yeni bir yöntem', Doktora Tezi, İnönü Üniversitesi Fen Bilimleri Enstitüsü, Malatya, Türkiye, 2021.
- [23] R. Uehara, S. Toda, and T. Nagoya, 'Graph isomorphism completeness for chordal bipartite graphs and strongly chordal graphs', *Discrete Appl Math (1979)*, vol. 145, no. 3, 2005, doi: 10.1016/j.dam.2004.06.008.
- [24] A. Kosorukoff and D. L. Passmore, *Social network analysis: Theory and applications*. Passmore, D. L, 2011. [Online]. Available: <https://books.google.com.tr/books?id=LrAnswEACA> AJ
- [25] F. Boudin, 'A Comparison of Centrality Measures for Graph-Based Keyphrase Extraction', in *6th International Joint Conference on Natural Language Processing, IJCNLP 2013 - Proceedings of the Main Conference*, 2013.
- [26] Anonymous, 'DUC 2002 Guidelines'. Accessed: Feb. 13, 2023. [Online]. Available: <https://www-nlpir.nist.gov/projects/duc/guidelines/2002.html>
- [27] D. Greene and P. Cunningham, 'Practical solutions to the problem of diagonal dominance in kernel document clustering', in *Proceedings of the 23rd international conference on Machine learning*, 2006, pp. 377–384.
- [28] H. P. Luhn, 'The automatic creation of literature abstracts', *IBM J Res Dev*, vol. 2, no. 2, pp. 159–165, 1958.
- [29] T. K. Landauer, P. W. Foltz, and D. Laham, 'An introduction to latent semantic analysis', *Discourse Process*, vol. 25, no. 2–3, pp. 259–284, 1998.
- [30] T. K. Landauer and S. T. Dumais, 'A solution to Plato's problem: The latent semantic analysis theory of acquisition, induction, and representation of knowledge.', *Psychol Rev*, vol. 104, no. 2, p. 211, 1997.

- [31] R. Mihalcea, 'Language independent extractive summarization', *ACL-05 - 43rd Annual Meeting of the Association for Computational Linguistics, Proceedings of the Conference*, pp. 49–52, 2005, doi: 10.3115/1225753.1225766.
- [32] R. Mihalcea and P. Tarau, 'Textrank: Bringing order into text', in *Proceedings of the 2004 conference on empirical methods in natural language processing*, 2004, pp. 404–411.
- [33] L. Page and S. Brin, 'The anatomy of a large-scale hypertextual Web search engine', *Computer Networks and ISDN Systems*, vol. 30, no. 1–7, pp. 107–117, Apr. 1998, doi: 10.1016/S0169-7552(98)00110-X.
- [34] 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.
- [35] L. Vanderwende, H. Suzuki, C. Brockett, and A. Nenkova, 'Beyond SumBasic: Task-focused summarization with sentence simplification and lexical expansion', *Inf Process Manag*, vol. 43, no. 6, pp. 1606–1618, 2007.
- [36] A. Haghighi and L. Vanderwende, 'Exploring content models for multi-document summarization', in *Proceedings of human language technologies: The 2009 annual conference of the North American Chapter of the Association for Computational Linguistics*, 2009, pp. 362–370.
- [37] A. Joshi, E. Fidalgo, E. Alegre, and L. Fernández-Robles, 'DeepSumm: Exploiting topic models and sequence to sequence networks for extractive text summarization', *Expert Syst Appl*, vol. 211, p. 118442, 2023.
- [38] J. Cheng and M. Lapata, 'Neural summarization by extracting sentences and words', *arXiv preprint arXiv:1603.07252*, 2016.
- [39] 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.
- [40] D. Parveen and M. Strube, 'Integrating importance, non-redundancy and coherence in graph-based extractive summarization', in *Twenty-Fourth International Joint Conference on Artificial Intelligence*, 2015.
- [41] D. Parveen, H.-M. Rams, and M. Strube, 'Topical coherence for graph-based extractive summarization', in *Proceedings of the 2015 conference on empirical methods in natural language processing*, 2015, pp. 1949–1954.
- [42] M. Tomer and M. Kumar, 'Multi-document extractive text summarization based on firefly algorithm', *Journal of King Saud University-Computer and Information Sciences*, vol. 34, no. 8, pp. 6057–6065, 2022.
- [43] R. Abbasi-ghalehtaki, H. Khotanlou, and M. Esmailpour, 'Fuzzy evolutionary cellular learning automata model for text summarization', *Swarm Evol Comput*, vol. 30, pp. 11–26, 2016.
- [44] R. M. Alguliyev, R. M. Aliguliyev, N. R. Isazade, A. Abdi, and N. Idris, 'A model for text summarization', *International Journal of Intelligent Information Technologies (IJIT)*, vol. 13, no. 1, pp. 67–85, 2017.
- [45] M. Mohamed and M. Oussalah, 'SRL-ESA-TextSum: A text summarization approach based on semantic role labeling and explicit semantic analysis', *Inf Process Manag*, vol. 56, no. 4, pp. 1356–1372, 2019.