

Automated biomedical material summarization using dwarf mongoose optimization with graph attention networks

Muthukumarasamy Pakkirisamy Balasubramanian¹, Perumal Senthamil Selvan^{2*} 

¹Government Polytechnic College, Department of Computer Engineering, Keelapaluvur Ariyalur, 62170, Tamil Nadu, India.

²University College of Engineering, Bharathidasan Institute of Technology Campus Anna University, Department of Pharmaceutical Technology, Tiruchirappalli, 620024, Tamil Nadu, India.

e-mail: balasubramanian12@gmail.com, senthamil77@gmail.com

ABSTRACT

Biomedical information encompasses various documents, offering researchers and medical practitioners valuable insights into the latest developments, validating and developing novel hypotheses, performing experiments, and interpreting results. Resources such as web information, multimedia documents, clinical trials, and medical reports also provide immense data. The growing size of these textual sources makes managing and extracting data challenging. Over recent decades, several automatic approaches have been developed to address text file issues for knowledge discovery and information extraction. Automatic biomedical text summarization techniques have been extensively explored to help researchers and clinicians handle large volumes of data. This paper presents an Automated Biomedical Document Summarization using Dwarf Mongoose Optimization with Graph Attention Networks (ABDS-DMOGAN) technique. The ABDS-DMOGAN model preprocesses biomedical documents to prepare them for summarization. It uses the GAN model for summarizing content from various biomedical documents. To enhance the GAN model's summarization results, the DMO algorithm is applied for hyperparameter tuning. Extensive simulations demonstrate the improved performance of the ABDS-DMOGAN model, with outcomes showing its significant superiority over other existing deep learning approaches.

Keywords: Bio-concrete data; Text summarization; Deep learning; Dwarf Mongoose optimizer; Graph attention networks.

1. INTRODUCTION

The fast evolution of text information for industry, academia, and research from different information sources [1], some are scientific articles, news, journals, medical records, databases, and books, turn out to be important landmarks in the progression stage of various technologies intended in the extraction of meaningful data in a wide range of application fields. Over the past, the development of a broad set of textual data is evident, especially in the biomedical sector [2]. Reviewing an immense volume of texts regularly, Healthcare professionals undergo a problem because of the meaningful information and knowledge this data includes, like drug usage, treatment, reactions, and symptoms, among others [3]. Biomedical text mining is designed with a need to address such information requirements by developing techniques for data extraction and data retrieval for biomedicine. Automatic text summarization refers to a vital text mining application targeted to summarize crucial data within documents in concise and more easily usable texts [4].

Automatic text summarization (ATS) presents a productive solution to access the increasing clinical and scientific literature by summarizing the source records while retaining their informative content [5]. It becomes a prominent topic in the domain of information retrieval research, especially in the biomedical and medical fields. A lot of biomedical data is been accessible to doctors makes it a great difficult to locate accurate data quickly [6]; so, automatic summarization could afford highlights specifically to medical needs. Also, attaining precise coherent information extraction, and succinct from credible published bio-medical resources to build a simplified summarization serves a vital role in medical education [7], clinical decision-making, and educating patients. Meanwhile, automation of this process brings better opportunities for users for gaining the most critical points of essential medical knowledge without looking into an immense volume of text, saving time in searching [8]. It is advisable to precisely detect scientifically sound published studies and

summarize selective studies from a given type (e.g., prognosis and intervention) of the medical query [9]. Various techniques were designed with the need to address the summarization challenges. Such techniques mostly use machine learning (ML) or graph methods. Graph-based techniques modelled the text as a graph and summarized it by examining the nodes [10].

This paper presents an Automated Biomedical Document Summarization using Dwarf Mongoose Optimization with Graph Attention Networks (ABDS-DMOGAN) technique. The proposed ABDS-DMOGAN model preprocesses the biomedical documents at the preliminary level to make them compatible with the summarization process. In addition, the ABDS-DMOGAN model utilizes the GAN model for the summarization of the content from several biomedical documents. For increasing the summarization results of the GAN model, the DMO algorithm is applied for the hyperparameter tuning procedure. An extensive range of simulations is carried out to reveal the improvised performance of the ABDS-DMOGAN method.

2. RELATED WORKS

The region of extraction automatic document summarization utilizing the DL technique and execution to Konkani language is low resource language as there are restricted resources, like speakers, tools, experts, or data in Konkani. During the presented method, Facebook's fastText pre-training word embedding was utilized for obtaining a vector representation for texts. Afterwards, the deep MLP approach was utilized as a supervised binary classifier task for auto-generating summary utilizing the feature vector [11]. A self-supervised system for extraction text summarization for biomedical texts. This technique utilizes abstracts for determining one of the informative content; afterwards, creates a summary to train a classifier method. The sentences in the abstract and texts are primary embedding employing Bidirectional Encoder Representations from Transformers (BERT). The authors utilized LR as our classifier method and utilized the feature of sentences embedded in classification [12].

A new summarizer technique which employs contextualized embedded created by the BERT technique, a DL system that modern established existing outcomes in various NLP tasks. The authors integrate various types of BERT including a clustering system for recognizing the very important and informative sentences of input documents [13]. A new word-embedded-based biomedical document summarization. Real dense vectors represent the biomedical words. The sentences can demonstrate by summing up the word vector which contains. The authors utilized a pre-training Word2vec approach of word vectors created in an integration of PMC, PubMed, and existing English Wikipedia dump texts [14].

A supervised extraction summarization system depending on conditional generative adversarial networks (CGAN) utilizing CNNs. Different preceding approaches that frequently utilize greedy approaches for selecting sentences, the authors utilize a novel technique to select sentences. Furthermore, the authors offer a network for biomedical word embedding that enhances the summarizer. An important role of the work has been establishing a novel loss function for the discriminator, constructing the discriminator to carry out optimum [15]. A novel biomedical document summarizer approach which incorporates 2 famous data mining approaches clustering and frequent item set mining. The K-means technique was exploited for clustering the same sentences. Next, the Apriori technique was executed for discovering the frequent item sets between the cluster sentences. At last, the salient sentences in every cluster are chosen for creating the summary utilizing the exposed frequent item sets [16]. A new summarization approach termed Clustering and Item set mining based Biomedical Summarization (CIBS). The summarization extractive biomedical models in the input texts and utilizes an item set mining technique for discovering essential topics [17].

3. THE PROPOSED MODEL

In this paper, a new ABDS-DMOGAN approach was introduced for the automated summarization of biomedical documents. The presented ABDS-DMOGAN technique follows a three-stage process: data pre-processing, summarization, and parameter tuning. Figure 1 represents the workflow of the ABDS-DMOGAN algorithm.

3.1. Data preprocessing

In the first stage, the presented ABDS-DMOGAN method preprocesses the biomedical documents. The redundant parts comprise the keywords, title, author's data, figures and tables, abstract, bibliography section, and headers of sections and subsections. It could be regarded that parts are redundant, as they could not execute under the approach summary that was employed to estimate the outline count. The removal stage was customized depending on the framework of input text and user preferences. Once it could be selected to contain sections and subsection titles in the summary, additional data was saved with every sentence for pointing out the section

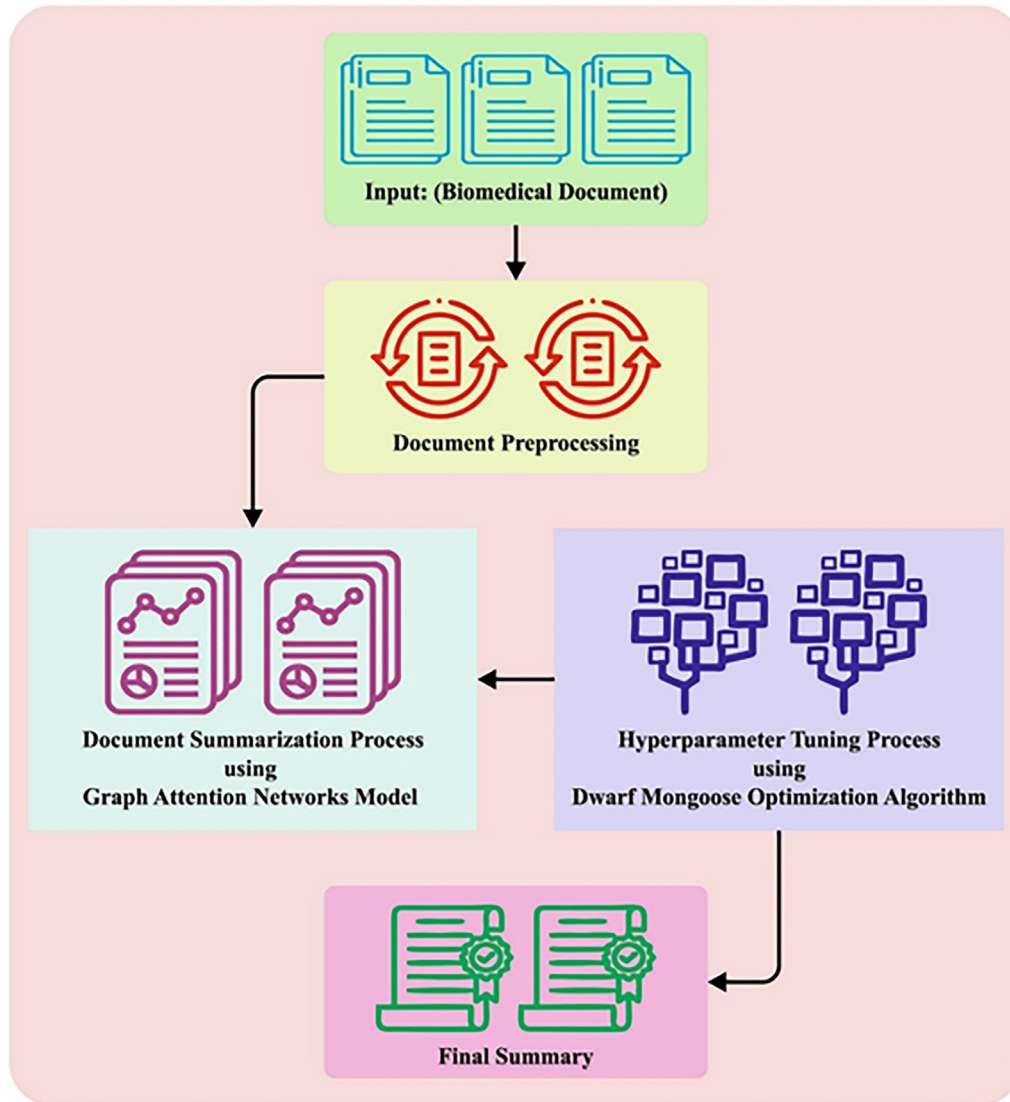


Figure 1: Workflow of ABDS-DMOGAN approach.

of text that sentence goes to. While the input of extraction features scripts of BERT can be text files whereas all the sentences carry out in various lines, and every sentence is tokenized, the pre-processing step rests with splitting an input text as different sentences and the tokenization step.

3.2. Document summarization using GAN

In the second stage, the biomedical documents data can be summarized using the GAN model. Assume a sentence $s = [w_1, w_2, \dots, w_p, \dots, w_n]$ with n length and w_i aspect target, we first map all the words into lower dimensional word embedding vector [18]. For all the words w_p , we get a single vector $x_i \in R^d$ where d refers to the dimension of word embedding space. Using an off-the-shelf dependency parser, we convert the sentence into a dependency graph. Every node characterizes the word and is related to the embedding vector as a local feature vector. The undirected edge between two words implies these two words are related syntactically.

The steps behind the proposed framework are

Initialization: The population of mongoose agents is initialized randomly in the search space.

Fitness Evaluation: The fitness of each agent is calculated based on the defined objective function.

Selection: The alpha mongoose (agent with the best fitness) is selected, and its position is noted.

Update: The positions of the other mongoose agents are updated based on the alpha's position, with some randomness introduced to explore the search space.

Convergence Check: The algorithm checks if the stopping criteria (e.g., maximum iterations or minimum fitness threshold) are met. If not, the process repeats.

Summarization: The optimized parameters are used in the GAN model for document summarization.

For the aspect objective with more than one word, we first replace the entire target word sequence with the special symbol “--target--”, and later passes the adapted sentence as a dependency parser. A graph attention network (GAT) is a kind of GNN and is a fundamental component. It proliferates features from the aspect syntax context to the aspect nodes. Assume a dependency graph with N nodes, whereas all the nodes are related to the local word embedding vector x , a single GAT layer calculates node representation by gathering neighborhood hidden state. Especially, an i -th node with hidden state h_l^i at l -th layer and the node neighbour $n[i]$ along with the hidden state, a GAT upgrades the hidden state node at $l + 1$ layer using multi-head attention. The updating procedure can be formulated as follows:

$$h_{l+1}^i = \parallel_{k=1}^K \sigma \left(\sum_{jk \in n[i]} \alpha_{lk}^{ij} W_{lk} h_l^j \right) \quad (1)$$

$$\alpha_{lk}^{ij} = \frac{\exp(f(a_{lk}^T [W_{lk} h_l^i \parallel W_{lk} h_l^j]))}{\sum_{u \in n[i]} \exp(f(a_{lk}^T [W_{lk} h_l^i \parallel W_{lk} h_l^u]))} \quad (2)$$

The Graph Attention Network (GAT) updates node features based on the attention scores between connected nodes. Here, a_{ij} represents the attention score between nodes i and j , and h_j is the feature vector of node j . The significance of a_{ij} lies in its ability to weigh the influence of neighboring nodes, allowing the GAT to focus on the most relevant connections when updating node features. This selective attention mechanism improves the model’s performance in capturing important relationships within the graph.

Where $\parallel$ denotes the vector concatenation, α_{lk}^{ij} represents the attention co-efficient of node i to its neighbour j in attention head k at the l layer $W_{lk} \in R^{\frac{D}{K} \times D}$ represent a linear conversion matrix for the input state. D represents a dimension of the hidden state. σ indicates the sigmoid function. f represents the LeakyReLU nonlinear function. During training, $a_{lk} \in R^{\frac{2D}{K}}$ denotes the attention context vector learned. We could write such feature propagation method as follows

$$H_{l+1} = GAT(H_l, A; \Theta_l) \quad (3)$$

In Equation 3, $H_l \in R^{N \times D}$ denotes the stacked state for every node at the l layer, and $A \in R^{N \times N}$ indicates the graph adjacent matrix. Θ_l represents the parameter set of GAT at the l layer.

3.3. Parameter optimization using DMO algorithm

For enhancing the summarization results of the GAN model, the DMO process is applied for the hyperparameter tuning procedure. The design of the DMO technique is based on the natural phenomena of the dwarf mongoose, deliberated as follows [19]. The author models each subgroup recognized by the dwarf mongoose population with the incorporation of the scouts, babysitter, and alpha (female and male) subgroups. The animal’s adaptive nature in their foraging predation was inspired and employed in the design stage. Firstly, the scouting group moves out to food sources but individuals allocated for babysitting are allowed to stay in the mound. It can be considered that the process of the lookout for food, called foraging, illustrates the exploration stage of the optimizer technique. Furthermore, the position of a novel food source enables a group of the population to settle down from the mound, the DMO model that as the intensification or exploitation stage.

Parameter Tuning: The DMO algorithm tunes the parameters by iteratively adjusting the positions of mongoose agents in the search space. Each agent’s position corresponds to a set of parameter values, and the fitness function evaluates the performance of these parameters. The criteria for selecting parameters include minimizing error rates and maximizing summarization accuracy. We have conducted preliminary experiments to determine the range of parameter values, ensuring that the DMO algorithm operates within an optimal range for the given task.

The representation of a set of populations that also signifies the whole population, is modelled through Equation 4. Meanwhile, the group population is frequently encompassed of the alpha, juvenile, and scout (involving babysitter) groups. Let PtP_t represent the population at iteration ttt. The population consists of NNN mongoose agents, where each agent's position is denoted by $X_i^t A = \pi r^2$ for $i = 1, 2, \dots, N$, $N_i = 1, 2, \dots, N$. Here, N is the total number of agents, and X_i^t represents the position of the iii-th mongoose in the search space at iteration ttt. As the role of alpha female (α) from the population, it can be employed in Equation 5 to calculate this special subgroup of alphas. It can be provided by assessing the fitness function (FF) of the whole group and also that fitness to know what individual is considered appropriate for alpha females (α).

$$X = \begin{bmatrix} x_{1,1} & x_{1,2} & \cdots & x_{1,d-1} & x_{1,d} \\ x_{2,1} & x_{2,2} & \cdots & x_{2,d-1} & x_{2,d} \\ \vdots & \vdots & & \vdots & \vdots \\ x_{n,1} & x_{n,2} & \cdots & x_{n,d-1} & x_{n,d} \end{bmatrix}, \quad (4)$$

$$\alpha = \frac{f \text{ it}_i}{\sum_{i=1}^n f \text{ it}_i}. \quad (5)$$

Likewise, the scout group represents the workforce of dwarf mongoose, which is extracted and computed in the whole population. To accomplish this, it can be proposed Equation 3 to characterize this derivation procedure, which simultaneously allows for the foraging activity and search for the newest sleeping mound.

$$X_{i+1} = \begin{cases} X_i - CF * phi * rand * [X_i - \vec{M}] & \text{if } \varphi_{i+1} > \varphi_i \\ X_i + CF * phi * rand * [X_i - \vec{M}] & \text{else} \end{cases} \quad (6)$$

In Equation 6, *rand* denotes a randomly generated value within [0,1], $CF = \left(1 - \frac{iter}{\text{Max}_{iter}}\right)^{\left(2 - \frac{iter}{\text{Max}_{iter}}\right)}$ shows the collective-volatile movement control variables and $\vec{M} = \sum_{i=1}^n \frac{X_i \times sm_i}{X_i}$ defines the mongoose movement to the newest *sm*, and $\varphi = \frac{\sum_{i=1}^n sm_i}{n}$.

The foraging and predatory actions within the mound of dwarf mongoose need that individual moves from one position to the other. This frequently needs an updated model to be employed to compute every position of the group member. They leverage the vocalization peep facility of alpha female (α) along with the randomly generated parameter *phi*, within [-1,1] to calculate the existing location of the individual. Figure 2 illustrates the flowchart of the DMO algorithm.

$$X_{i+1} = X_i + phi * peep. \quad (7)$$

Consider a scenario where the position of a mongoose is given by $X_i^t = [0.5, 0.3]$ at iteration ttt. The alpha mongoose has a position $X_\alpha = [0.7, 0.6]$. The updated position for the next iteration might be calculated as:

$$X_i^{t+1} = X_i^t + \alpha (X_\alpha - X_i^t) + \beta \cdot \text{random_vector},$$

where α is a learning rate, and β introduces randomness. This results in the mongoose moving closer to the alpha's position while exploring the search space.

Another crucial milestone considered by foraging and predatory behaviors is the discovery of the newest sleeping mound. This newest mound was considered for changing under the iteration procedure of the optimizer. To model the find of *sm*, Equation 8 is employed to compute this and enables to average φ the *sm* value.

$$sm_j = \frac{fit_{i+1} - fit_i}{\max\{|fit_{i+1}, fit_i\}}. \quad (8)$$

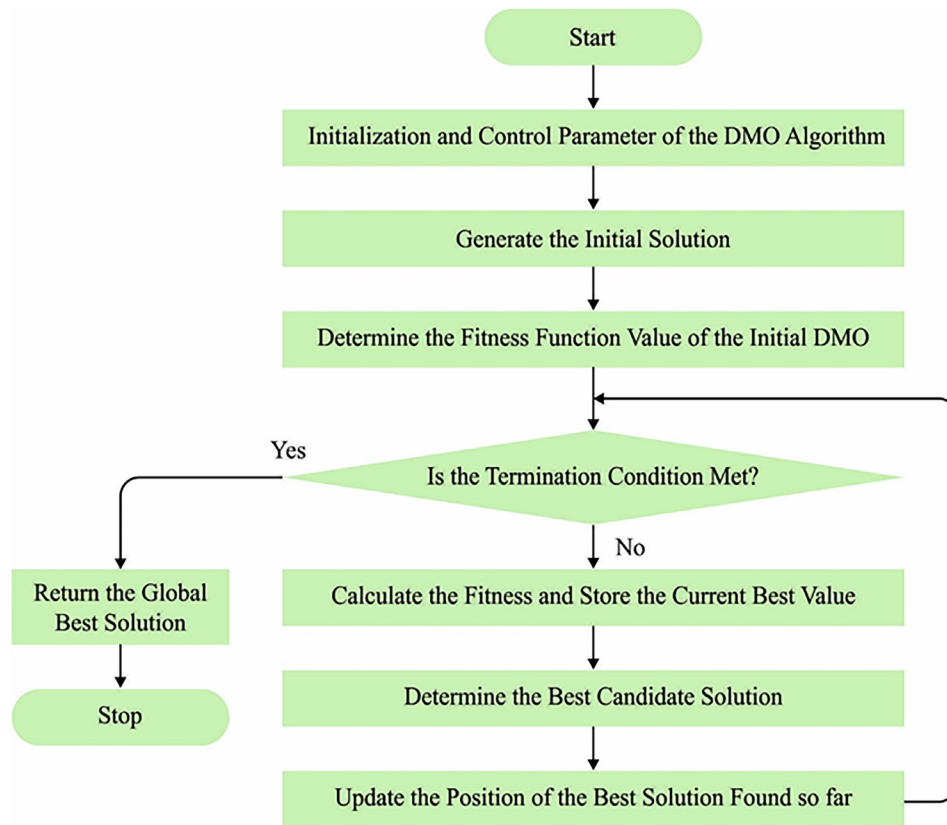


Figure 2: Flowchart of DMO algorithm.

In this section, the process for the application of the mathematical model defined is encoded as follows:

- a. Initialize each control parameter
- b. The group of population is generated initially
- c. Subgroup the whole population into babysitters, alpha (female and male), and scouts
- d. Define the presented amount of searching agents by deducting the babysitter from the whole population
- e. Fix the exchange rate for babysitting tasks as L
- f. While the ending criteria are unfulfilled, do the following:
 - i. Calculate the fitness of the mongoose group or population
 - ii. Activate and set a time counter
 - iii. Employ Equation 5 for deducing the size of alpha female
 - iv. Enhanced DMO for Constraint Engineering Design Problem
 - v. Upgrade the location of an effective food source based on Equation 7
 - vi. Iterate over all the individuals and calculate the fitness of X_i
 - vii. Derive the SM for the population based on Equation 8
 - viii. Define the movement vector $\vec{M}$
 - ix. Exchange the babysitter
 - x. Calculate the scout group based on Equation 6
 - xi. Upgrade the solution so far
- g. Return the optimum solution

The above mentioned process was transformed into pseudocode and later performed as a DMO technique.

Algorithm 1: Pseudocode of DMO Algorithm

begin

Initializing the algorithm parameters:

[peep]

Initializing the mongoose populations (searching agents): n

Initializing the babysitters count: bs

Set $n = n - bs$

Set babysitter interchange parameter L

For $iter - 1$: max iter

 Estimate the fitness of mongoose

 Set time counter C

 Define the alpha dependent upon

$$\alpha = \frac{fit_i}{\sum_{i=1}^n fit_i}$$

 Create a candidate food position

$$X_{i+1} = X_i + phi * peep$$

 Estimate novel fitness of

X_{i+1} Estimate sleeping mound

$$sm_i = \frac{fit_{i+1} - fit_i}{\max\{|fit_{i+1}, fit_i|\}}$$

 Calculate the average value of the sleeping mound found.

$$\varphi = \frac{\sum_{i=1}^n sm_i}{n}$$

 Calculate the movement vector utilizing

$$\vec{M} = \sum \frac{X_i \times sm_i}{X_i}$$

 Interchange babysitters if $C \geq L$, and set

 Initializing bs position and computing fitness

$$fit_i \leq \alpha$$

 Put on the scout mongoose's next position.

$$X_{i+1} = \begin{cases} X_i - CF * rand * [X_i - \vec{M}] & \text{if } \varphi_{i+1} > \varphi_i & \text{Exploration} \\ X_i - CF * rand * [X_i - \vec{M}] & \text{else} & \text{Exploration} \end{cases}$$

The DMO method has derived a FF to achieve enhanced classifier outcomes. It has determined a positive value for indicating the higher outcome of the candidate solutions. Here, the diminished classifier error rate is the FF, as offered in Equation 9. The fitness function in Equation 9 plays a critical role in guiding the DMO algorithm towards optimal solutions. It balances exploration and exploitation by rewarding mongoose agents for positions that lead to better summarization accuracy.

$$fitness(x_i) = Classifier \ Error \ Rate(x_i)$$

$$= \frac{no. \ of \ misclassified \ instances}{Total \ no. \ of \ instances} * 100$$

4. EXPERIMENTAL VALIDATION

The experimental result analysis of the ABDS-DMOGAN method is tested on the PubMed dataset [20], which encompasses the data samples in the json format. The abstract, sections, and body are every sentence tokenized. The json objects encompass several variables namely: section_names, article_id, article_text, sections, and abstract_text.

In Table 1, the overall summarization results of the ABDS-DMOGAN model with the GAN model are stated clearly. Figure 3 shows the overall assessment of the ABDS-DMOGAN model in terms of ROUGE-1. The results demonstrated that the ABDS-DMOGAN model reports enhanced ROUGE-1 values over several iterations. For instance, with iteration 1, the ABDS-DMOGAN model attains ROUGE-1 of 76.17%.

Next, with iteration 4, the ABDS-DMOGAN algorithm attains ROUGE-1 of 76.14%. Meanwhile, with iteration 8, the ABDS-DMOGAN method attains ROUGE-1 of 78.36%. Moreover, with iteration 10, the ABDS-DMOGAN approach attains ROUGE-1 of 77.59%.

Figure 4 displays the overall assessment of the ABDS-DMOGAN method in terms of ROUGE-2. The figure shows that the ABDS-DMOGAN technique reports enhanced ROUGE-2 values over several iterations. For example, with iteration 1, the ABDS-DMOGAN approach achieves ROUGE-2 of 36.85%. Next, with iteration 4, the ABDS-DMOGAN method attains ROUGE-2 of 35.37%. In the meantime, with iteration 8, the ABDS-DMOGAN approach gains ROUGE-2 of 35.37%. Likewise, with iteration 10, the ABDS-DMOGAN technique reaches ROUGE-2 of 36.56%.

Figure 5 displays the overall assessment of the GAN method in terms of ROUGE-1. The figure exhibited that the GAN model reports enhanced ROUGE-1 values over several iterations. For instance, with iteration 1,

Table 1: Overall summarization results of ABDS-DMOGAN and GAN methods with varying iterations.

NO. OF ITERATIONS	ABDS-DMOGAN		GAN MODEL	
	ROUGE1	ROUGE2	ROUGE1	ROUGE2
1	76.17	36.85	74.88	34.37
2	75.37	35.99	74.43	35.34
3	75.56	36.6	74.29	34.33
4	76.14	35.37	75.42	34.74
5	75.18	35.01	74.09	35.33
6	78.65	36.77	74.37	35.70
7	79.37	36.04	77.74	35.33
8	78.36	35.37	77.98	35.85
9	78.34	36.92	76.20	34.02
10	77.59	36.56	75.54	35.90

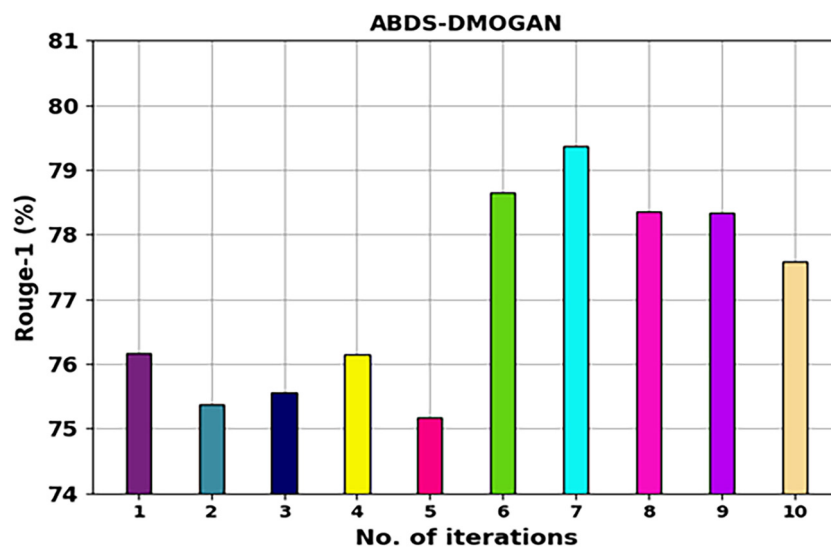


Figure 3: Rouge-1 analysis of ABDS-DMOGAN method under varying iterations.

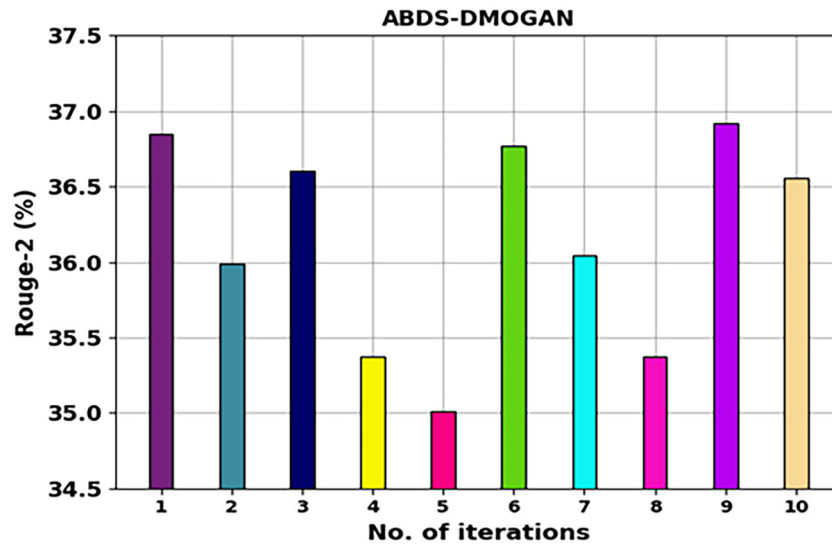


Figure 4: Rouge-2 analysis of ABDS-DMOGAN method under varying iterations.

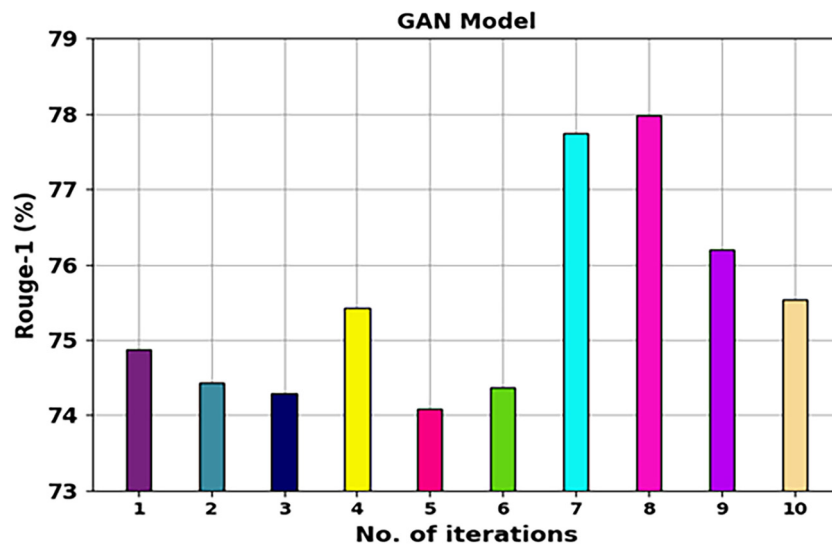


Figure 5: Rouge-1 analysis of the GAN method under varying iterations.

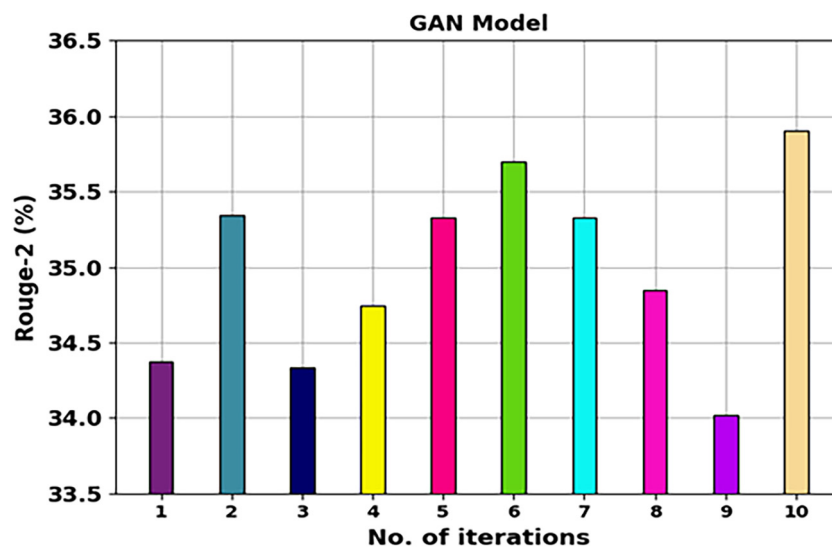


Figure 6: Rouge-2 analysis of the GAN method under varying iterations.

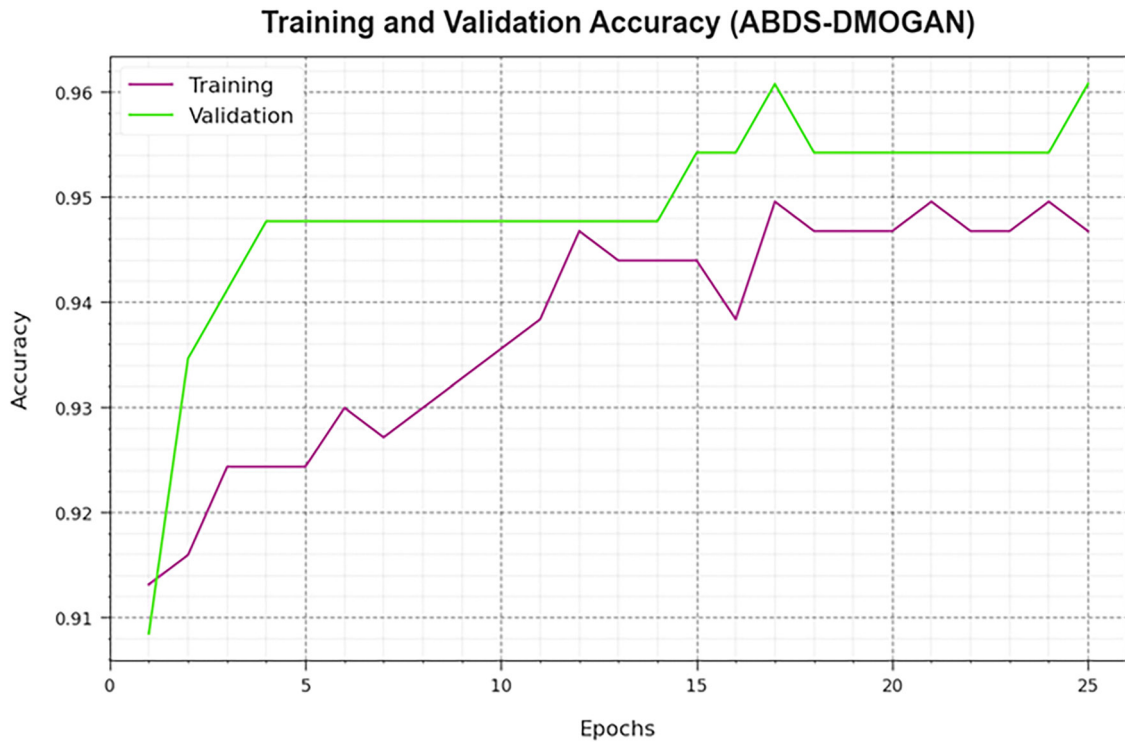


Figure 7: Accuracy curve of the ABDS-DMOGAN methodology.

the GAN method attains ROUGE-1 of 74.88%. Then, with iteration 4, the GAN model attains ROUGE-1 of 75.42%. Meanwhile, with iteration 8, the GAN model attains ROUGE-1 of 77.98%. Further, with iteration 10, the GAN model attains ROUGE-1 of 75.54%.

Figure 6 shows the overall assessment of the GAN model in terms of ROUGE-2. The results validated that the GAN model reports enhanced ROUGE-2 values over several iterations. For example, with iteration 1, the GAN model attains ROUGE-2 of 34.37%. Next, with iteration 4, the GAN model attains ROUGE-2 of 34.74%. In the meantime, with iteration 8, the GAN model attains ROUGE-2 of 35.85%. Besides, with iteration 10, the GAN model attains ROUGE-2 of 35.90%.

Figure 7 inspects the accuracy of the ABDS-DMOGAN technique during the training and validation process on the test dataset. The figure specifies that the ABDS-DMOGAN technique reaches increasing accuracy values over increasing epochs. In addition, the increasing validation accuracy over training accuracy reveals that the ABDS-DMOGAN algorithm learns efficiently on the test dataset.

The loss analysis of the ABDS-DMOGAN approach at the time of training and validation is demonstrated on the test dataset in Figure 8. The figure indicates that the ABDS-DMOGAN technique reaches closer values of training and validation loss. The ABDS-DMOGAN technique learns efficiently on the test dataset.

Figure 9 examines the accuracy of the GAN technique during the training and validation process on the test dataset. The figure notifies that the GAN technique reaches increasing accuracy values over increasing epochs. In addition, the increasing validation accuracy over training accuracy exhibits that the GAN technique learns efficiently on the test dataset.

The loss analysis of the GAN method at the time of training and validation is demonstrated on the test dataset in Figure 10. The figure specifies that the GAN approach reaches closer values of training and validation loss. The GAN technique learns efficiently on the test dataset.

To display the improved performance of the ABDS-DMOGAN approach, a comparative ROUGE analysis is made in Table 2 [21, 22]. Figure 11 demonstrates a brief ROUGE-1 inspection of the ABDS-DMOGAN method with other existing models. The results indicate the effectual outcomes of the ABDS-DMOGAN model with an increasing ROUGE-1 value of 79.37%. At the same time, the GAN, DL-ALSTM, LSTM, BERT-large, BioBERT-pubmed, Bayesian BS, and BERT-base models obtain reducing ROUGE-1 values of 77.98%, 76.86%, 75.28%, 75.04%, 73.76%, 72.88%, and 72.57% respectively.

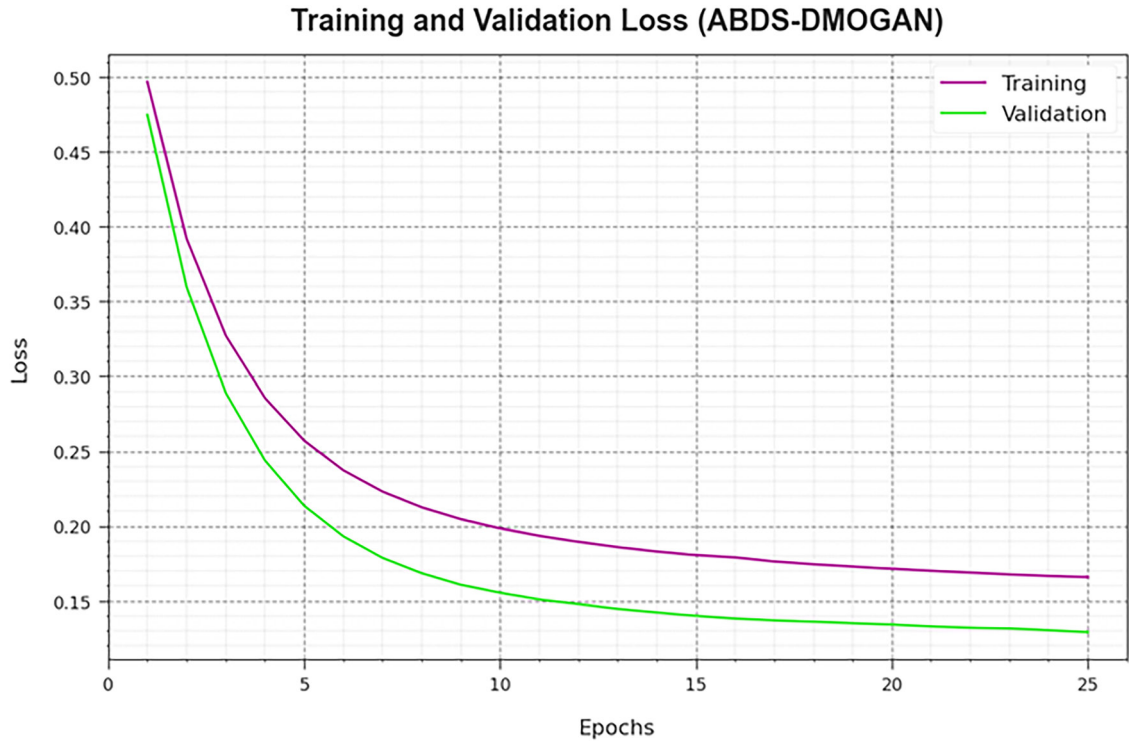


Figure 8: Loss curve of the ABDS-DMOGAN methodology.

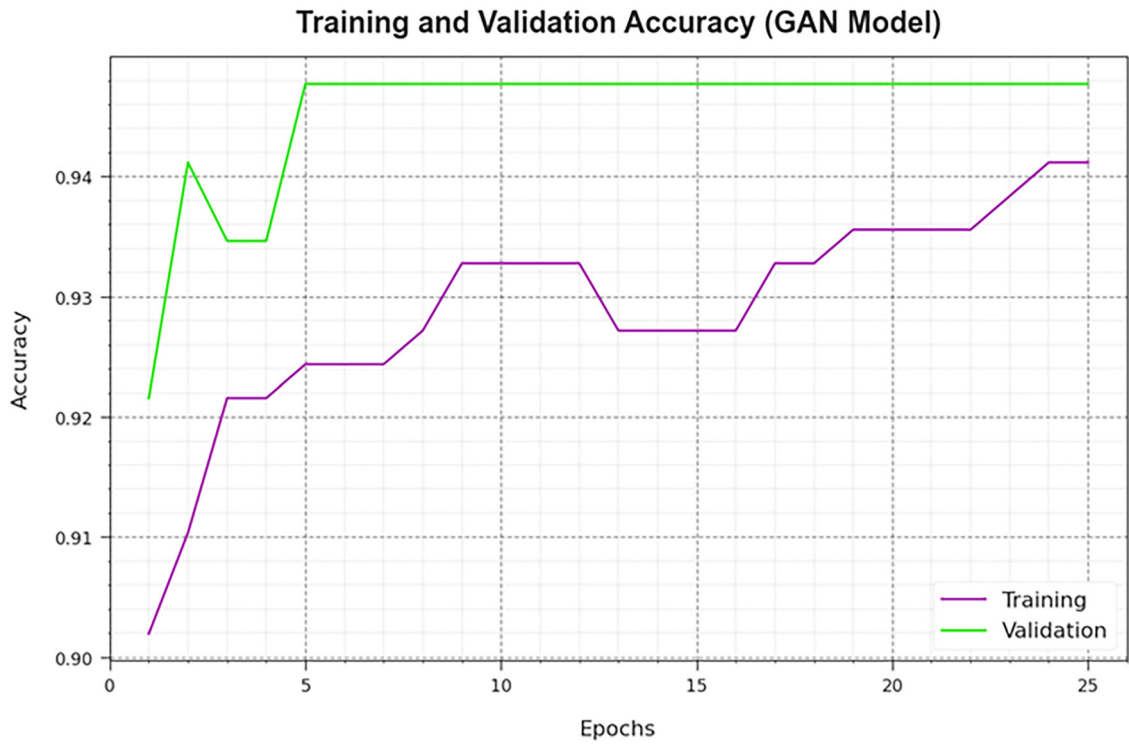


Figure 9: Accuracy curve of the GAN methodology.

The proposed method differs from traditional clustering and frequent itemset mining techniques in that it utilizes the DMO algorithm to optimize summarization parameters, leading to more accurate and contextually relevant summaries [23]. While clustering methods group similar documents, our approach focuses on optimizing the summarization process itself. The frequent itemset mining techniques are effective for identifying common patterns but may not capture the nuanced relationships that the DMO-optimized GAN model can achieve.

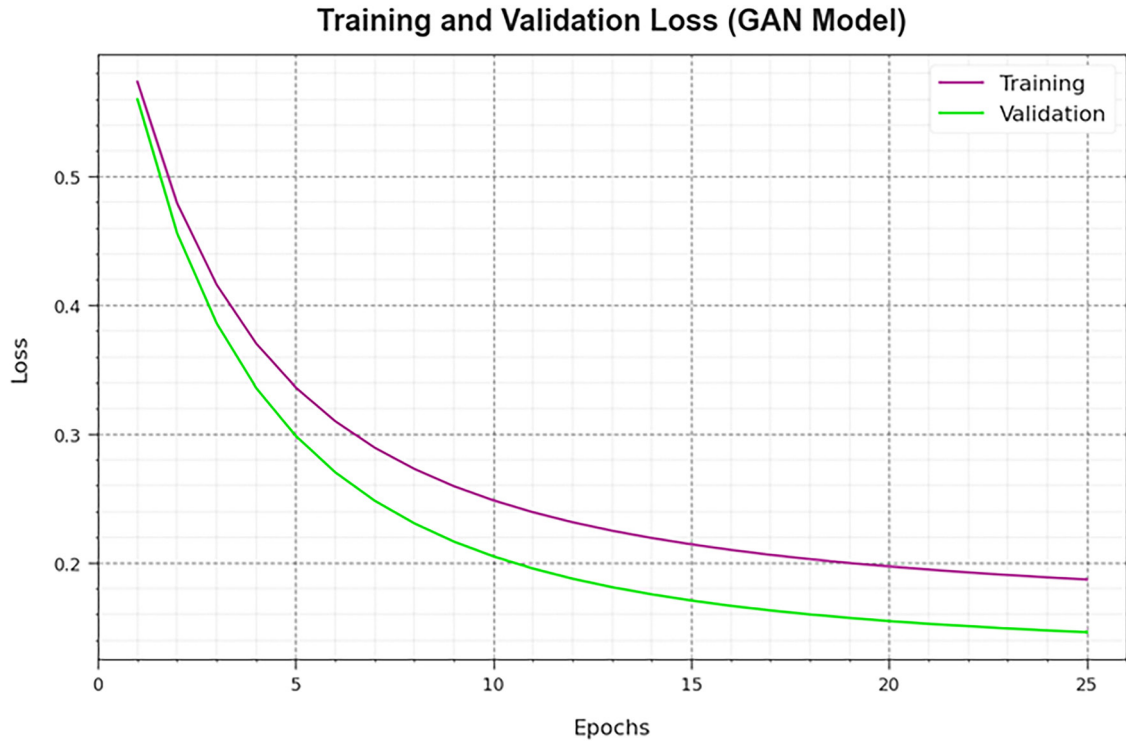


Figure 10: Loss curve of the GAN methodology.

Table 2: Comparative outcome of ABDS-DMOGAN approach with other techniques.

METHODS	ROUGE-1 (%)	ROUGE-2 (%)
ABDS-DMOGAN	79.37	36.04
GAN Model	77.98	35.85
DL-ALSTM	76.86	35.12
LSTM	75.28	35.62
BERT-large	75.04	33.12
BioBERT-pubmed	73.76	32.03
Bayesian BS	72.88	31.43
BERT-base	72.57	31.10

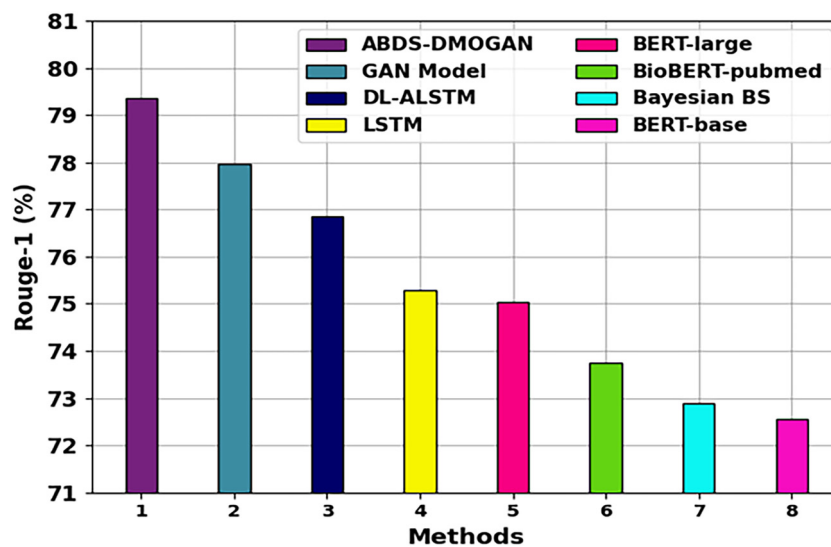


Figure 11: Rouge-1 analysis of ABDS-DMOGAN approach with other techniques.

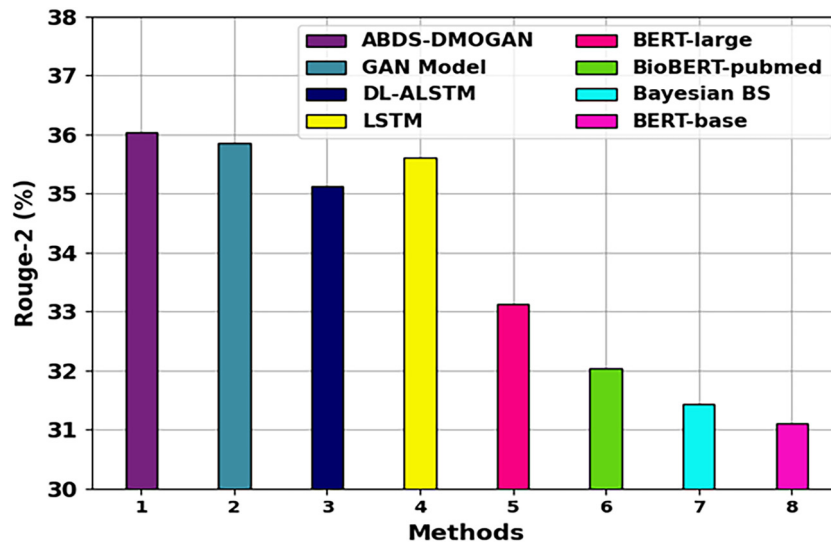


Figure 12: Rouge-2 analysis of ABDS-DMOGAN approach with other techniques.

Figure 12 exhibits a comparative ROUGE-2 inspection of the ABDS-DMOGAN technique with other existing models. The results indicate the effectual outcomes of the ABDS-DMOGAN algorithm with an increasing ROUGE-1 value of 36.04% [24]. Simultaneously, the GAN, DL-ALSTM, LSTM, BERT-large, BioBERT-pubmed, Bayesian BS, and BERT-base models obtain reducing ROUGE-1 values of 35.85%, 35.12%, 35.62%, 33.12%, 32.03%, 31.43%, and 31.10% respectively.

These results confirmed the improved summarization performance of the ABDS-DMOGAN model compared to other summarization approaches.

5. CONCLUSION

In this paper, we have developed a new ABDS-DMOGAN approach for the automated summarization of biomedical documents. The presented ABDS-DMOGAN technique follows a three-stage process: data pre-processing, summarization, and parameter tuning. In the first stage, the proposed ABDS-DMOGAN model preprocesses the biomedical documents at the preliminary level to make them compatible with the summarization process. Followed by, the ABDS-DMOGAN model utilizes the GAN model for the summarization of the content from several biomedical documents. For enhancing the summarization results of the GAN model, the DMO algorithm is applied for the hyperparameter tuning procedure. To demonstrate the improvised performance of the ABDS-DMOGAN method, a wide range of simulations were performed. The simulation results report the significant performance of the ABDS-DMOGAN model over other existing DL models.

6. BIBLIOGRAPHY

- [1] RAI, A., SANGWAN, S., GOEL, T., *et al.* "Query specific focused summarization of biomedical journal articles", In: *Proceedings of the 16th Conference on Computer Science and Intelligence Systems (FedCSIS)*, pp. 91–100, 2021. doi: <http://doi.org/10.15439/2021F128>.
- [2] GERO, Z., HO, J.C. "Uncertainty-based self-training for biomedical keyphrase extraction", In: *Proceedings of the 2021 IEEE EMBS International Conference on Biomedical and Health Informatics (BHI)*, pp. 1–4, 2021.
- [3] WANG, M., WANG, M., YU, F., *et al.*, "A systematic review of automatic text summarization for biomedical literature and EHRs", *Journal of the American Medical Informatics Association*, v. 28, n. 10, pp. 2287–2297, 2021. doi: <http://doi.org/10.1093/jamia/ocab143>. PubMed PMID: 34338801.
- [4] MISHRA, R., BIAN, J., FISZMAN, M., *et al.*, "Text summarization in the biomedical domain: a systematic review of recent research", *Journal of Biomedical Informatics*, v. 52, pp. 457–467, 2014. doi: <http://doi.org/10.1016/j.jbi.2014.06.009>. PubMed PMID: 25016293.
- [5] KADHAR, S.A., GOPAL, E., SIVAKUMAR, V., *et al.*, "Optimizing flow, strength, and durability in high-strength self-compacting and self-curing concrete utilizing lightweight aggregates", *Matéria*, v. 29, n. 1, e20230336, 2024. doi: <http://doi.org/10.1590/1517-7076-rmat-2023-0336>.

- [6] LI, Z., LIN, H., SHEN, C., *et al.* “Cross2Self-attentive bidirectional recurrent neural network with BERT for biomedical semantic text similarity”, In: *Proceedings of the 2020 IEEE International Conference on Bioinformatics and Biomedicine (BIBM)*, pp. 1051–1054, Seoul, Korea (South), 2020. doi: <http://doi.org/10.1109/BIBM49941.2020.9313452>.
- [7] VAN VEEN, D., VAN UDEN, C., BLANKEMEIER, L., *et al.* “Clinical text summarization: adapting large language models can outperform human experts”, *Research Square*, 2023. In press. doi: <http://doi.org/10.21203/rs.3.rs-3483777/v1>.
- [8] PARTHASAARATHI, R., BALASUNDARAM, N., NAVEEN ARASU, A., “Analysing the impact and investigating Coconut Shell Fiber Reinforced Concrete (CSFRC) under varied loading conditions”, *Journal of Advanced Research in Applied Sciences and Engineering Technology*, v. 35, n. 1, pp. 106–120, 2024.
- [9] CANDEMIR, S., ANTANI, S., XUE, Z., *et al.*, “Novel method for storyboarding biomedical videos for medical informatics”, In: *Proceedings of the 2017 IEEE 30th International Symposium on Computer-Based Medical Systems (CBMS)*, pp. 127–132, Thessaloniki, Greece, 2017. doi: <http://doi.org/10.1109/CBMS.2017.57>.
- [10] D’SILVA, J., SHARMA, U., “Automatic text summarization of konkani texts using pre-trained word embeddings and deep learning”, *Iranian Journal of Electrical and Computer Engineering*, v. 12, n. 2, pp. 1990, 2022. doi: <http://doi.org/10.11591/ijece.v12i2.pp1990-2000>.
- [11] SRINIVASAN, S.S., MUTHUSAMY, N., ANBARASU, N.A., “The structural performance of fiber-reinforced concrete beams with nanosilica”, *Matéria*, v. 29, n. 3, e20240194, 2024. doi: <http://doi.org/10.1590/1517-7076-rmat-2024-0194>.
- [12] MORADI, M., DORFFNER, G., SAMWALD, M., “Deep contextualized embeddings for quantifying the informative content in biomedical text summarization”, *Computer Methods and Programs in Biomedicine*, v. 184, pp. 105117, 2020. doi: <http://doi.org/10.1016/j.cmpb.2019.105117>. PubMed PMID: 31627150.
- [13] ROUANE, O., BELHADEF, H., BOUAKKAZ, M., “Word embedding-based biomedical text summarization”, In: *Emerging Trends in Intelligent Computing and Informatics: Data Science, Intelligent Information Systems and Smart Computing 4*, pp. 288–297, 2020. doi: http://doi.org/10.1007/978-3-030-33582-3_28.
- [14] NAVEEN ARASU, A., NATARAJAN, M., BALASUNDARAM, N., *et al.*, “Utilizing recycled nanomaterials as a partial replacement for cement to create high performance concrete”, *Global NEST Journal*, v. 6, n. 25, pp. 89–92, 2023.
- [15] MORAVVEJ, S.V., MIRZAEI, A., SAFAYANI, M., “Biomedical text summarization using conditional generative adversarial network (CGAN)”, *arXiv*, 2023. In press.
- [16] ROUANE, O., BELHADEF, H., BOUAKKAZ, M., “Combine clustering and frequent itemsets mining to enhance biomedical text summarization”, *Expert Systems with Applications*, v. 135, pp. 362–373, 2019. doi: <http://doi.org/10.1016/j.eswa.2019.06.002>.
- [17] NAVEEN ARASU, A., NATARAJAN, M., BALASUNDARAM, N., *et al.*, “Optimization of high performance concrete by using nano materials”, *Research on Engineering Structures Materials*, v. 3, n. 9, pp. 843–859, 2023.
- [18] NAVEEN ARASU, A., RANJINI, D., PRABHU, R., “Investigation on partial replacement of cement by GGBS”, *Journal of Critical Reviews*, v. 7, n. 17, pp. 3827–3831, 2020.
- [19] ABACHA, A.B., M’rabet, Y., ZHANG, Y., *et al.*, “Overview of the MEDIQA 2021 shared task on summarization in the medical domain”, In: *Proceedings of the 20th Workshop on Biomedical Language Processing*, pp. 74–85, 2021. doi: <http://doi.org/10.18653/v1/2021.bionlp-1.8>.
- [20] NAVEEN ARASU, A., NATARAJAN, M., NAVEEN ARASU, A., “A comprehensive microstructural analysis for enhancing concrete’s longevity and environmental sustainability”, *Journal of Environmental Nanotechnology*, v. 13, n. 2, pp. 368–376, 2024. doi: <http://doi.org/10.13074/jent.2024.06.242584>.
- [21] AGUSHAKA, J.O., EZUGWU, A.E., OLAIDE, O.N., *et al.*, “Improved dwarf mongoose optimization for constrained engineering design problems”, *Journal of Bionics Engineering*, v. 20, n. 3, pp. 1263–1295, 2023. doi: <http://doi.org/10.1007/s42235-022-00316-8>. PubMed PMID: 36530517.
- [22] NAVEEN ARASU, A., NATARAJAN, M., BALASUNDARAM, N., *et al.*, “Development of high performance concrete by using nano material graphene oxide in partial replacement of cement”, *AIP Conference Proceedings*, v. 2861, pp. 050008, 2023. doi: <http://doi.org/10.1063/5.0158487>.

- [23] ALMASOUD, A.S., HASSINE, S., AL-WESABI, F., *et al.*, “Automated multidocument biomedical text summarization using deep learning model”, *Computers, Materials and Continua*, v. 71, n. 3, pp. 5799–5815, 2022.
- [24] ALTMAMI, N.I., MOHAMED, E.B.M., “Automatic summarization of scientific articles: a survey”, *Journal of King Saud University-Computer and Information Sciences*, v. 34, n. 4, pp. 1011–1028, 2022. doi: <http://doi.org/10.1016/j.jksuci.2020.04.020>.