

[https://doi.org/10.52326/jes.utm.2025.32\(1\).03](https://doi.org/10.52326/jes.utm.2025.32(1).03)
UDC 004.032.26:519.7:81'322=966.8



A GRAPH CONVOLUTIONAL NETWORK AND RECURRENT NEURAL NETWORK ENSEMBLE FOR EXTRACTIVE TEXT SUMMARISATION IN THE HAUSA LANGUAGE

Anas Shehu *, ORCID: 0000-0002-5307-1457,
Musa Karatu, ORCID: 0000-0002-6337-3117,
Abubakar Roko, ORCID: 0000-0002-3489-5346,
Sirajo Abdullahi, ORCID: 0000-0002-4346-621X

Federal University Birnin Kebbi, 860222, Nigeria

* Corresponding authors: Anas Shehu, 2200264009@pg.fubk.edu.ng

Received: 01. 29. 2025

Accepted: 03. 04. 2025

Abstract. Automatic Text Summarisation (ATS) is crucial for managing information overload, especially in low-resource languages like Hausa. This study proposes a hybrid extractive approach that combines Graph Convolutional Networks (GCN) and Recurrent Neural Networks (RNN) to improve sentence ranking accuracy. By integrating GCN's structural learning with RNN's sequential modeling, the method overcomes limitations of existing graph-based techniques. The research was conducted in Visual Studio IDE using Python 3.12.4, with key libraries like NLTK, Pandas, and NetworkX. A pre-processed Hausa news dataset was tokenised, normalised, and vectorised using TF-IDF and a pre-trained Hausa FastText model to build a sentence similarity graph. GCNs propagated sentence embeddings, while RNNs refined rankings by capturing sequential dependencies. Experiments on 113 Hausa news articles showed the GCN-RNN model outperformed Modified PageRank, achieving higher ROUGE-1 precision (90.00) and balanced F1-scores. The Wilcoxon Signed-Rank Test confirmed significant improvements. Despite added computational overhead, the approach remains feasible for moderate datasets, with scalability as a key future focus. This study offers a robust and contextually coherent approach to Hausa text summarisation, advancing extractive summarisation techniques and multilingual ATS research. Future work will focus on optimising model efficiency and scalability while exploring transformer-based architectures for further enhancements.

Keywords: *Natural Language Processing ~ Hausa Language, Extractive Text Summarisation, Deep Learning Ensembles, Automatic Text Summarisation.*

Rezumat. Sumarizarea automată a textului (ATS) este crucială pentru gestionarea supraîncărcării cu informații, în special în limbile cu resurse reduse, cum ar fi hausa. Acest studiu propune o abordare extractivă hibridă care combină rețelele convoluționale grafice (GCN) și rețelele neuronale recurente (RNN) pentru a îmbunătăți acuratețea clasificării propozițiilor. Prin integrarea învățării structurale a GCN cu modelarea secvențială a RNN, metoda depășește limitele tehnicilor existente bazate pe grafuri. Cercetarea a fost efectuată

în Visual Studio IDE folosind Python 3.12.4, cu biblioteci cheie precum NLTK, Pandas și NetworkX. Un set de date de știri hausa preprocesat a fost tokenizat, normalizat și vectorizat folosind TF-IDF și un model hausa FastText pre-antrenat pentru a construi un grafic de similaritate a propozițiilor. GCN-urile au propagat încorporările de propoziții, în timp ce RNN-urile au rafinat clasamentele prin captarea dependențelor secvențiale. Experimentele pe 113 articole de știri hausa au arătat că modelul GCN-RNN a depășit performanța Modified PageRank, atingând o precizie ROUGE-1 mai mare (90,00) și scoruri F1 echilibrate. Testul Wilcoxon Signed-Rank a confirmat îmbunătățiri semnificative. În ciuda costurilor de calcul suplimentare, abordarea rămâne fezabilă pentru seturi de date moderate, scalabilitatea fiind un obiectiv cheie în viitor. Acest studiu oferă o abordare robustă și coerentă din punct de vedere contextual a rezumării textului în limba hausa, avansând tehnicile de rezumare extractivă și cercetarea ATS multilingvă. Lucrările viitoare se vor concentra pe optimizarea eficienței și scalabilității modelului, explorând în același timp arhitecturi bazate pe transformatoare pentru îmbunătățiri suplimentare.

Cuvinte cheie: *prelucrarea limbajului natural, limba hausa, sumarizare extractivă a textului, ansambluri de învățare profundă, sumarizare automată a textului.*

1. Introduction

The rapid growth of digital information has made Automatic Text Summarisation (ATS) an essential tool for information retrieval and text classification, helping to mitigate information overload [1]. Extractive summarisation, which selects key sentences from a document to generate concise summaries, has been successfully applied across various languages, including Marathi [2], Indonesian [3], and Urdu [4]. However, achieving accurate sentence ranking in Hausa text summarisation remains a challenge due to limited contextual understanding and graph construction constraints [5].

Extractive text summarisation serves as a promising method for text summarisation that works by selecting a subset of existing words, phrases, or sentences in the original text to form the summary [6]. Despite the effectiveness of graph-based methods in extractive text summarisation, accurately ranking sentences in Hausa language remains a challenge, leading to suboptimal summaries [5]. Integrating a deep learning algorithm holds promise for improving summarisation accuracy within these methods [7].

However, Bichi, Samsudin, Hassan, Hasan and Rogo [5], employed a modified PageRank algorithm to rank sentences for extractive summarisation in Hausa text. The modified algorithm considers the relationships between sentences in the text to determine their importance, which is a departure from the original PageRank algorithm. The modified PageRank algorithm faces challenges such as its reliance on graph construction contextual, limited understanding, scalability issues, and insufficient semantic awareness, all of which can lead to less effective summaries. To address these limitations, this study proposes a hybrid extractive summarisation model that integrates Graph Convolutional Networks (GCNs) and Recurrent Neural Networks (RNNs). GCNs effectively capture semantic relationships, while RNNs refine sentence rankings by leveraging sequential dependencies, improving summarisation accuracy.

This research contributes to multilingual ATS by offering a robust, contextually aware extractive summarisation method for Hausa. Future work will focus on scalability improvements and integration with transformer-based models to enhance performance on larger datasets.

2. Related Work

2.1 Graph-Based Approaches

Graph theory has been a foundational method for automatic text summarization, providing a structured way to rank sentences based on their relationships. LexRank, introduced in 2004, was one of the earliest graph-based summarization models, using eigenvector centrality to compute sentence importance [8].

This approach has since inspired a wide range of research, with adaptations for various domains, input types, and languages.

Belwal, *et al.* [9], proposed a graph-based extractive summarization model that integrates topic modeling to mitigate redundancy and enhance semantic coherence. Similarly, Jalil, *et al.* [10] developed Grapharizer, a multi-document summarization (MDS) technique, addressing grammar inconsistencies, missing key information, and redundancy. Their approach improved summarization accuracy by 23.05% on ROUGE metrics compared to baseline models.

For Arabic text summarization, AL-Khassawneh and Hanandeh [11] introduced a textual graph-based approach, outperforming state-of-the-art models like Latent Semantic Analysis (LSA) [12], and Lexical Cohesion and Entailment-based segmentation (LCEAS) [13]. Evaluations on the Essex Arabic Summary Corpus (EASC) showed enhanced precision and informativeness.

In the context of low-resource languages, Bichi, Samsudin, Hassan, Hasan and Rogo [5] introduced a graph-based ranking method for Hausa extractive summarization, leveraging normalized common bigrams count to outperform TextRank, LexRank, Centroid-based, and BM25-TextRank in ROUGE-1, ROUGE-2, and ROUGE-L evaluations. However, this method lacks a deep learning component, limiting its ability to capture contextual and sequential dependencies.

Further studies have explored graph-based sentence scoring for summarization in diverse languages. Bhattacharya, *et al.* [14], introduced a sentence-centric graph-based text mining model, demonstrating its effectiveness in crowdsourcing applications. Similarly, Sarwadnya and Sonawane [2] developed a graph-based summarizer for Marathi, though it lacked semantic ranking. Ghos, *et al.* [15], applied a rule-based extractive graph technique for Bangla news, outperforming previous sentence extraction models.

Recent advancements in graph neural networks (GNNs) have further improved summarization performance. Cui, *et al.* [16], introduced Topic-Aware Graph Neural Networks (Topic-GraphSum) to capture inter-sentence relationships in long documents. Their approach outperformed traditional methods like SumBasic, LexRank, and LSA on datasets such as CNN/Daily Mail and The New York Times (NYT).

2.2 Algorithm Approach

Beyond graph-based methods, algorithmic models-ranging from machine learning to evolutionary algorithms - have contributed significantly to extractive summarization.

Verma and Nidh [17], introduced a deep learning-based approach that selects key sentences while preserving content meaning and coherence, combining Restricted Boltzmann Machines and Fuzzy Logic. Alami *et al.* [18], improved unsupervised neural network-based summarization by integrating Word2Vec embeddings and ensemble learning, demonstrating superior performance over Bag-of-Words (BoW) representations.

Hybrid techniques have also been explored. Gulati, *et al.* [19], proposed a BM25+ and TextRank hybrid model, showing that combining graph-based and probabilistic methods

improves summarization across news, research papers, and blogs. Girsang and Amadeus [3], developed an Ant System Algorithm for Indonesian text summarization, demonstrating its effectiveness in content synthesis and readability.

In evolutionary algorithms, Mojrian and Mirroshande [20] introduced Multi-Document Text Summarization using a Quantum-Inspired Genetic Algorithm (MTSQIGA). This model outperformed state-of-the-art summarization techniques, including Centroid-based methods [21], Probabilistic Latent Semantic Analysis (PLSA) [22], SRSUM [23]. However, fine-tuning hyperparameters remains a challenge, potentially leading to redundancy and suboptimal summaries.

For Turkish, Akülker [24], explored TF-IDF and PageRank, confirming their effectiveness in extractive summarization and suggesting further testing on larger datasets with multiple thresholds.

2.3 Advancements in Neural Networks for Summarization

While graph-based and algorithmic methods remain prominent, recent advancements in deep learning – particularly RNNs and GCNs – have shown promise in extractive summarization.

The Multi-GraS model [25] demonstrated how Bi-LSTMs (a type of RNN) and GCNs can jointly capture syntactic and semantic relationships, improving extractive summarization for English datasets (CNN/Daily Mail). This highlights the potential of neural graph-based hybrid models.

Hausa text summarization, however, remains in early research stages, with most methods relying solely on graph-based PageRank variations [5]. To advance Hausa summarization, this study proposes an integrated GCN-RNN approach, leveraging GCNs for structural relationships and RNNs for sequential dependencies. This integration aims to enhance sentence ranking accuracy, ensuring contextually rich and semantically coherent summaries.

While graph-based and algorithmic approaches have achieved success in extractive text summarization, Hausa text summarization remains underexplored. Existing studies primarily employ modified PageRank, lacking deep learning integration for semantic and sequential understanding.

To address the lack of integration deep learning in Hausa text summarisation, this study integrates GCNs and RNNs to enhance sentence ranking accuracy for Hausa extractive summarisation.

By leveraging GCNs for sentence relationships and RNNs for sequential dependencies, this hybrid approach aims to bridge the gap between structural coherence and contextual depth.

3. Methodology

The proposed method utilized a hybrid framework that integrate the strength of both GCN and RNN as shown in Figure 1 to enhance the accuracy of sentence ranking in Hausa for graph based extractive text Summarisation.

The GCN methods provide a structured representation of the text, capturing semantic relationships between sentences. While RNN algorithms learned the intricate patterns and representations from data, potentially to improve the accuracy of sentence ranking.

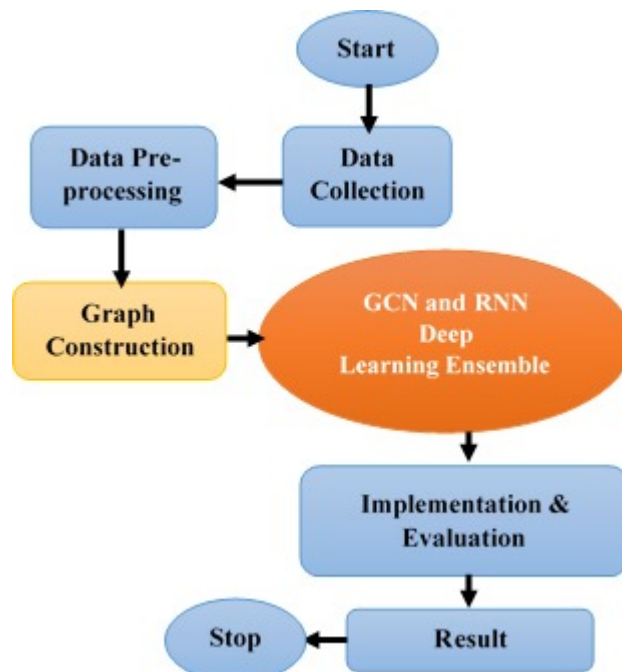


Figure 1. Proposed Deep Learning Ensemble Framework.

3.1 Dataset

The experiments was conducted on a Hausa summarisation evaluation dataset comprising five (5) Hausa publishers such as AMINIYYA, HAUSA LEADERSHIP NEWSPAPER, BBC HAUSA, RFI HAUSA, and VOA, HAUSA, which gives a total of 113 Hausa news articles adopted from [5], the Table 1 below describes the dataset used in evaluation the proposed method.

Table 1

Summary of the 113 Hausa news articles adopted from [5]

Source	Number of articles	Total Words	Number of manual summaries	Summary length, %	Average number of sentences in the article	Average number of sentences in the summaries
Aminiyya newspaper	25	7859	50	20	10	2
Hausa Leadership newspaper	21	9043	42	20	11	2
BBC Hausa	23	6314	46	20	9	2
RFI Hausa	23	6800	46	20	5	1
VOA Hausa	21	6509	42	20	11	2

The dataset was normalised using Pandas one of the Python libraries that support cleaning and standardising of text data, such as lowercasing, removing special characters, and handling whitespace in the dataset. The sentences were tokenised using the Natural Language Toolkit (NLTK) and the Hausa stopword was removed from the dataset.

3.2 Text Vectorisation and Graph Construction

A pre-trained Hausa Fast text Model was used for word representation learning (word embeddings) developed by Facebook which offers pre-trained word vectors for 157

languages, and Hausa Language is inclusive. The TF-IDF vectoriser was used in transforming the article and reference summary into vectors and calculating the cosine similarity between these vectors, the sentence similarity scores were computed based on the vectorised text from TF-IDF, and a similarity graph to rank the sentences based on scores was created using GCN.

3.3 Sentence Ranking using GCNs

GCNs are increasingly used in **extractive text summarisation** to **rank sentences** based on their relevance and importance [26]. The sentence ranking was achieved after the text has been represented as a node and edge as the relation between the nodes.

Mathematically, the resulting graph can be represented as:

Graph representation:

- The document is represented as a graph:

$$G = (V, E), \quad (1)$$

where: V is the set of nodes (sentences).

E is the set of edges (relationships between sentences).

- The adjacency matrix A_{ij} is defined such that A_{ij} = weight of the edge between sentence i and sentence j .

Node Feature Representation

- Each sentence (node) is associated with an **initial feature matrix** $H^{(0)}$, where: $H^{(0)} \in \mathbb{R}^{N \times d}$, with N being the number of sentences and d being the feature dimension.

Each row of $H^{(0)}$ represents the feature vector of a sentence.

Graph Convolution Operation

- The node embeddings are updated at each layer l using the GCN propagation rule:

$$H^{(l+1)} = \sigma(\tilde{D}^{-1/2} \tilde{A} \tilde{D}^{-1/2} H^{(l)} W^{(l)}), \quad (2)$$

where: $\tilde{A} = A + I$ is the adjacency matrix with self-loops (identity matrix I added).

$\tilde{D}$ is the diagonal degree matrix of $\tilde{A}$.

$W^{(l)}$ is the trainable weight matrix at layer l .

σ is the non-linear Rectified Linear Unit (ReLU) activation function.

$H^{(l)}$ is the feature matrix at layer l .

The GCN learned better sentence representations based on the propagating information across the graph structure when the sentence representations iteratively update through graph convolutional layers. Each layer propagates information between neighboring nodes (sentences), and the simple scoring function to assign a score to each sentence based on its learned representation was used by GCN to rank the sentences based on their importance, and a higher score indicates more important or central sentences. The simple scoring function mathematically is represented as $H^{(l)}$:

$$Score_i = f(H_i^{(l)}), \quad (3)$$

where: f is a non-linear function.

3.4 Evaluation metrics

Recall-Oriented Understudy for Gisting Evaluation (Rouge) as a widely used set of metrics for evaluating the quality of text summarisation, the study used Rouge-1, Rouge-2, and Rouge-L metrics to provide a good indication of the informativeness and fluency of the generated summary based on the following research objectives:

- Rouge-1 (Unigram Overlap) was used to evaluate the impact of combining multiple features on the quality of summaries generated for Hausa news articles such as word-level, sentence-level, and graph-level features.
- Rouge-2 (Bigram Overlap) was also used in monitoring sentence scores to avoid multiple sentences conveying similar information.
- Rouge-L helps in evaluating the overall semantic and syntactic similarity between the generated summary and the reference summary.

4. Proposed Method

The Graph convolutional network and Recurrent neural network are a class of deep learning models that operate on graph-structured data and handle sequential data respectively, the idea is to integrate the ability of both of them to improve the precision of sentence ranking on Hausa text summarisation. This method was divided into two tasks:

- Graph convolutional network task.
- And Recurrent neural network task.

4.1 The Graph convolutional network task (GCN)

The GCN learns and handles better sentence representations after been represented as nodes based on the propagating information across the graph structure, and the sentence representations are iteratively updated through graph convolutional layers. Each layer propagates information between neighboring nodes (sentences), the core idea is to have a better message-passing mechanism, where each node aggregates and computes new representations based on the features of its neighboring nodes and the edges connecting them to know the relationships and dependencies within the graph and produce sequence encoding to the RNN as input.

4.2 The Recurrent neural network task

RNNs are a specialised type of neural network architecture that are designed to handle sequential data, such as text, speech, or time series information, unlike the traditional feedforward neural networks, where each input is processed independently, RNNs have connections that form directed cycles. This method, used this structure of RNN to maintain memory and a hidden state to capture information about previous inputs, enabling them to handle sequential data effectively supplied by GCN. Mathematically it is represented as:

$$h_t = f(W_h h_{t-1} + W_x x_t + b_h), \quad (4)$$

where: W_h is the weight matrix for the hidden state.

W_x is the weight matrix for the input.

b_h is the bias vector.

f is an activation function.

x_t is an input to the RNN.

h_t is the hidden state.

h_{t-1} previous hiding state.

The evaluation result indicates higher precision values of 90.00, 87.01 and 85.00 in Rouge-1 than in Rouge-2 and L respectively. It shows a higher significant matching of words, very good coherence, and achieved semantic and syntactic similarity between the generated summary and the reference summary when compared with GCN only. The proposed method shows a **balanced performance** across the three Rouge metrics, with relatively high recall values of 37.50, 35.32, and 35.42, this indicates the method captures not only key terms but also contextual and structural aspects of the summaries compared to previous results of GCN.

The F1-score values for the proposed method **52.94 Rouge-1, 50.24 Rouge-2, and 50.00 Rouge-L indicate the method** performs well in identifying key terms from the reference summaries, balancing precision (relevance of selected words) and recall (coverage of reference words). Also, the method achieves a good balance between precision and recall for capturing contextual relationships between words. This result supports the idea that the summary maintains coherence and preserves meaningful phrases. The structural similarity with both recall and precision were effectively balanced, ensuring the generated summaries align well with the reference summaries in terms of sentence structure and flow.

These demonstrate that the proposed method strikes a good balance between precision and recall, the results suggest a **robust and effective summarisation method**, particularly when compared to recall-only results, as F1-scores consider both the quality and completeness of the generated summaries.

5. Result and Discussion

The performance of the Modified PageRank method and the GCN-RNN method was evaluated using the Rouge metrics of Rouge-1, Rouge-2, and Rouge-L. These metrics assess the quality of text summarisation by measuring unigram overlap, bigram overlap, and overall structural similarity between the generated and reference summaries. The precision, recall, and F1-score values for both methods are compared below.

Table 2

Comparison of the proposed method with other Methods

Evaluation Metrics	Performance Metrics	GCN	Modified PageRank	Proposed Method
Rouge-1	Precision	70.89	73.73	90.00
	Recall	18.06	70.90	37.50
	F1-score	28.79	72.28	52.94
Rouge-2	Precision	34.26	32.91	87.01
	Recall	8.72	35.94	35.32
	F1-score	13.90	34.36	50.24
Rouge-L	Precision	60.76	70.31	85.00
	Recall	15.48	70.65	35.42
	F1-score	24.68	70.48	50.00

Note: GCN - Graph Convolutional Networks.

Table 2 presents the results of the experiments using the proposed method, GCN, and the Modified PageRank. The Modified PageRank has a - Precision of 73.73, Recall 70.90, and F1-score 72.28 in Rouge-1 which demonstrates higher coherence (recall). The proposed approach GCN- RNN, with a significant difference of 33.4% and a balanced F1-score rate

compared to the Modified PageRank – Precision 90.00, Recall 37.50, and F1-score 52.94. The result of GCN-RNN is indicative of its ability to capture a broader range of key terms. However, its lower precision compared to Modified PageRank, suggests that proposed method may include less relevant terms. Also, the GCN-RNN method exhibits significantly higher precision with about 16.27% relative to the Modified PageRank, highlighting its effectiveness in selecting highly relevant matching of words, although at the cost of a lower recall value.

The F1-score for GCN-RNN suggests a reasonable balance, but the trade-off between precision and recall is clearer. The Modified PageRank method shows moderate performance on Rouge-2, with relatively balanced precision and recall values of 32.91, and 35.94, respectively. Again, the GCN-RNN method outperforms Modified PageRank in precision by a significant rate of 54%, indicating better coherence and selection of relevant bigrams. Despite the slightly lower recall, the GCN-RNN's higher F1-score reflects its superior capability to balance relevance and coverage for bigram matches about a difference of 16%.

Moreover, in the Structural similarity, the Modified PageRank method achieves highly balanced performance across Precision, Recall, and F1-score values of – 70.31, 70.65, and 70.48, respectively for ROUGE-L, indicating Modified PageRank ability to preserve sentence structure and flow. In contrast, the GCN-RNN method prioritises precision, achieving a significantly higher value in ROUGE-L metric with 14.6% difference. The lower recall and moderate F1-score of the Modified PageRank suggests that while the Modified PageRank captures structural aspects effectively, and the overall performance of the proposed method is shown below.

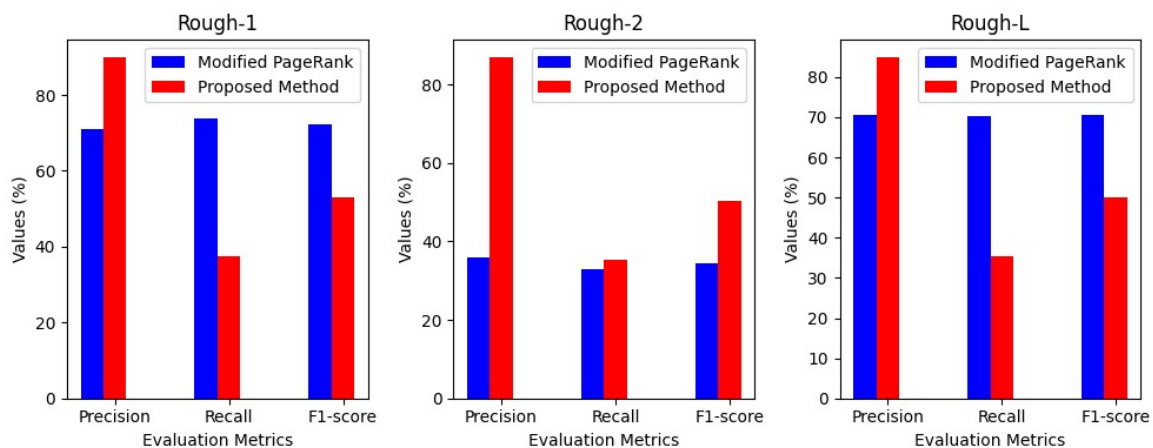


Figure 4. Overall performance of the proposed method.

The performance of both methods based on Precision as shown in Figure 4 above, indicates that the GCN-RNN method outperforms the Modified PageRank method across all ROUGE metrics in terms of precision, demonstrating its effectiveness in selecting relevant and high-quality content in the Generated Summary. In the Recall performance, the Modified PageRank method consistently achieves higher recall values, reflecting its broader coverage of the reference summaries. Also, the performance of the F1-Score shows The GCN-RNN method achieves higher F1-scores in ROUGE-2 and ROUGE-L, indicating better overall performance in capturing contextual and structural relationships.

5.1 Analysis for Statistical Significance

The **Wilcoxon statistic** is a **non-parametric** statistical test used to compare **paired or independent samples** when the data does not necessarily follow a normal distribution [27].

It is commonly used in machine learning and statistical analysis to compare the performance of two algorithms. The **Wilcoxon Signed-Rank Test was used in this study, and the performance of Modified PageRank and GCN-RNN methods was paired.** The Wilcoxon test results indicate a statistically significant difference between the precision scores of Modified PageRank and GCN-RNN methods (p-value: 1.91e-06). This suggests that GCN-RNN significantly outperforms Modified PageRank in terms of precision, as the differences are not due to random chance in all the Rouge 1, 2, and L respectively.

The Wilcoxon test results indicate a statistically significant difference between the F1-scores of Modified PageRank and GCN-RNN (p-value: 1.91e-06). This suggests that Modified PageRank significantly outperforms GCN-RNN in terms of F1-score in Rouge 1 and L only, implying better balance between precision and recall in this evaluation.

However, in Rouge 2 the Wilcoxon test results indicate a statistically significant difference between the F1-scores of Modified PageRank and GCN-RNN (p-value: 1.91e-06). This confirms that GCN-RNN significantly outperforms Modified PageRank in terms of F1-score, suggesting it is a more effective model. Also in terms of Recall the result show a statistically significant difference between the recall scores of Modified PageRank and GCN-RNN (p-value: 0.024) all the Rouge 2. This confirms that GCN-RNN significantly outperforms Modified PageRank in terms of recall, demonstrating its superior ability to retrieve relevant results.

5.2 A Time Complexity Analysis of the GCN and RNN Ensemble Model

This section presents the time complexity analysis of the GCN and RNN ensemble for Hausa text summarisation. The complexity of different components of the pipeline are analysed.

Preprocessing Steps. The preprocessing step sets the stage for effective sentence ranking and graph-based learning. By meticulously converting text into meaningful numerical formats and capturing inter-sentence relationships, it ensures that the subsequent models – such as GCNs and RNNs – receive high-quality inputs, thereby enhancing the quality and coherence of the generated summaries.

The text data is transformed using TF-IDF with a vocabulary size of d (max. features = 793):

- TF-IDF computation complexity: $O(Nd)$,

where N is the number of sentences (nodes), and d is the maximum number of features.

- Cosine similarity computation: $O(N^2d)$.

Since cosine similarity requires pairwise comparison between all N sentences, it results in a quadratic complexity.

Overall complexity of preprocessing (TF-IDF + Similarity Computation):

$$O(Nd) + O(N^2d) = O(N^2d) \quad (5)$$

Graph Construction. Converting the Similarity Matrix into a Graph each similarity matrix is converted into an adjacency matrix representation:

- Constructing the graph takes $O(N^2)$ operations as it directly maps the similarity matrix into an undirected graph.
- Extracting edges for the GCN (edge_index computation) involves iterating through nonzero elements, leading to $O(E)$ complexity, where E is the number of edges.

The Graph construction complexity:

$$O(N^2) + O(E) \quad (6)$$

The GCN consists of two convolutional layers. The time complexity per layer is:

$$(EF) + O(NFH) + O(NHO), \quad (7)$$

where: E is number of edges in the graph; F - feature dimension (793); H - hidden layer dimension (64); O - output feature dimension (64).

Since the GCN has two layers, the total complexity is:

$$O(2(EF + NFH + NHO)) = O(EF + NFH + NHO) \quad (8)$$

Since $E \approx O(N)$ for a sparse graph, we approximate this as:

$$O(NF + NFH + NHO) = O(NF + NH^2) \quad (9)$$

Recurrent Neural Network

The RNN (LSTM) processes the N sentences as a sequence. The per-step complexity is:

$$O(H^2 + HF) \quad (10)$$

Since the LSTM runs for N steps, the total complexity is:

$$O(N(H^2 + HF)) \quad (11)$$

Sentence Ranking

Sorting the N sentences based on RNN output has a complexity of $O(N \log N)$,

Overall Complexity of GCN-RNN

Summing up the key components:

$$O(N^2d) + O(N^2) + O(NF + NH^2) + O(N(H^2 + HF)) + O(N \log N) \quad (12)$$

Since $d = F$ (793), and assuming $F < NH$, we approximate the dominant terms:

$$O(N^2d) + O(N^2) + O(NF) + O(NH^2) + O(NHF) \quad (13)$$

For large N, the $O(N^2d)$ and $O(N^2)$ terms dominate, giving the final complexity:

$$O(N^2d) \quad (14)$$

This quadratic complexity primarily results from TF-IDF vectorisation and cosine similarity computations. The ensemble of GCN and RNN adds only a linear term $O(NH^2 + NHF)$, making it scalable for moderate-sized datasets like those used in this research ($N < 1000$), but costly for very large text corpora. The scalability consideration for larger datasets $> 10,000$ is the next direction of this research.

5.3 Time Complexity Analysis of the Modified PageRank

The Modified PageRank algorithm follows a graph-based ranking mechanism where nodes (sentences) are ranked based on their importance in summarisation. The computational complexity can be broken down into:

- **Graph Construction:** Constructing the similarity graph requires computing similarity scores for all sentence pairs, leading to an $O(N^2)$ complexity.
- **Power Iteration for PageRank:** PageRank relies on iterative computation where convergence is typically reached within I iterations. Each iteration involves a sparse matrix-vector multiplication, resulting in a per-iteration complexity of $O(E)$. Thus, for I iterations, the total complexity is:

$$O(IE) \approx O(IN) \quad (15)$$

Since we have a small, constant number of iterations I , the final complexity is near-linear:

$$O(N^2 + IN) \quad (16)$$

For large datasets, where $I < NI$, this approximates to $O(N^2)$, making it suitable for moderate-scale text summarisation.

GCN provides a scalable solution for structured text representations, Modified PageRank is effective for ranking-based summarisation, and the ensemble approach balances feature learning and sequence modelling but at a higher computational cost.

6. Conclusions

This research detailed a comprehensive study on Hausa text summarisation, integrating various natural language processing and machine learning techniques to enhance summary quality and model transparency. The experimental setup, conducted within a robust Visual Studio IDE environment using Python 3.12.4 and key libraries such as NLTK, Pandas, and NetworkX, ensured a solid foundation for reproducible and effective research.

The study began by implementing a meticulously pre-processed dataset from multiple Hausa news publishers, where text cleaning, tokenisation, and normalisation set the stage for further analysis. Text vectorisation was achieved using a pre-trained Hausa FastText model combined with TF-IDF and cosine similarity measures, leading to the construction of a sentence similarity graph. This graph served as the basis for sentence ranking via GCNs, which efficiently propagated and refined sentence representations.

However, the inherent limitations of GCNs in capturing multi-word expressions and complex structural dependencies necessitated the development of a hybrid approach. The proposed method, which integrates GCN with RNNs, leveraged the sequential modelling strengths of RNNs to complement the structural insights from the GCN. This ensemble method demonstrated superior performance in balancing precision and recall, as evidenced by improved Rouge metrics and statistically significant results via the Wilcoxon Signed-Rank Test.

The results and discussion section further revealed that while the Modified PageRank method achieved strong structural preservation, the GCN-RNN ensemble outperformed it in precision and overall F1-score, highlighting its ability to effectively identify key terms and maintain contextual coherence. The time complexity analysis indicated that although the ensemble method introduces additional computational overhead – primarily due to TF-IDF vectorisation and cosine similarity calculations – it remains feasible for moderate-sized datasets. Scalability considerations for larger corpora were identified as a promising direction for future research.

In conclusion, the integration of GCN and RNN in this study provides a robust and balanced approach to Hausa text summarisation, combining effective feature learning with advanced sequence modelling. The findings underscore the potential of this ensemble method to produce summaries that are both precise and contextually coherent, while also offering valuable insights into computational trade-offs and scalability challenges. Future work will focus on optimising the model for larger datasets and further refining the balance between performance and computational efficiency.

Conflicts of Interest: The authors declare no conflicts of interest.

References

1. Bilal, K.; Ali, S.Z.; Muhammad, U.; Inayat, K.; Badam, N. Exploring the landscape of automatic text summarization: a comprehensive survey. *IEEE Access* 2023, 11, 109819-109840.
2. Sarwadnya, V.V.; Sonawane, S.S. Marathi Extractive Text Summarizer using Graph Based Mode. In: *Proceedings of the Fourth International Conference on Computing Communication Control and Automation (ICCUBEA)*, Pune, India, 16-18 August 2018, pp. 1-6.
3. Girsang, A.S.; Amadeus, F.J. Extractive Text Summarization for Indonesian News Article Using Ant System Algorithm. *Journal of Advances in Information Technology* 2023, 14, 295-301.
4. Nawaz, A.; Bakhtyar, M.; Baber, J.; Ullah, I.; Noor, W.; Basi, A. Extractive Text Summarization Models for Urdu Language. *Information Processing and Management* 2020, 57, 1-12.
5. Bichi, A.A.; Samsudin, R.; Hassan, R.; Hasan, L.R.A.; Rogo, A.A. Graph-based extractive text summarization method for Hausa text. *PLOS ONE* 2023, 5, 1-15.
6. Ramesh, R.; Rajan, B. Extractive Text Summarization Using Graph Based Ranking Algorithm And Mean Shift Clustering. In: *Proceedings of the In Proceeding of International Conference on Recent Trends in Computing, Communication & Networking Technology(ICRTCCNT)* Chennai, Tamil Nadu, 18-19 October 2019, pp. 1-8.
7. Amirsina, T.; Rouzbeh-A, S.; Yaser, K.; Nader, T.; Edward-A, F. Natural language processing advancements by deep learning: A survey. *arXiv preprint arXiv:2003.01200* 2020, 2, 1-23.
8. Erkan, G.; Radev, D.R. LexRank: Graph-based Lexical Centrality as Saliency in Text Summarization. *Journal of Artificial Intelligence Research* 2004, 22, 457-479.
9. Belwal, R.C.; Rai, S.; Gupta, A. A new graph-based extractive text summarization using keywords or topic modeling. *Journal of Ambient Intelligence and Humanized Computing* 2020, 10, 1-16, doi:https://doi.org/10.1007/s12652-020-02591-x.
10. Jalil, Z.; Nasir, M.; Alazab, M.; Nasir, J.; Amjad, T.; Alqammaz, A. Grapharizer: A Graph-Based Technique for Extractive Multi-Document Summarization. *Electronics* 2023, 12, 1-26.
11. AL-Khassawneh, Y.A.; Hanandeh, E.S. Extractive Arabic Text Summarization-Graph-Based Approach. *Electronics* 2023, 12, 437.
12. El-Haj, M.; Kruschwitz, U.; Fox, C. Experimenting with automatic text summarisation for arabic. In: *Proceedings of the Human Language Technology. Challenges for Computer Science and Linguistics: 4th Language and Technology Conference*, LTC 2009, Poznan, Poland, November 6-8, 2009, Revised Selected Papers 4, Poznan, Poland, 2011, pp. 490-499.
13. Al-Khawaldeh, F.; Samawi, V. Lexical cohesion and entailment based segmentation for arabic text summarization (lceas). *WSCIT* 2015, 5, pp. 51-60.
14. Bhattacharya, P.; Verma, J.P.; Bhargav, S.; Bhavsar, M. GETS : Sentence Scoring Scheme in Graph-based Extractive Text Summarization for Text Mining Applications. *Information* 2023, 14, pp. 1-28.
15. Ghos, P.P.; Shahariar, R.; Khan, M.A.H. A Rule Based Extractive Text Summarization Technique for Bangla News Documents. *IJ. Modern Education and Computer Science* 2018, 12, pp. 44-53.
16. Cui, P.; Hu, L.; Liu, Y. Enhancing Extractive Text Summarization with Topic-Aware Graph Neural Networks. In: *Proceedings of the Proceedings of the 28th International Conference on Computational Linguistics*, Barcelona, Spain 8-13 December 2020, 2020, pp. 5360-5371.
17. Verma, S.; Nidh, V. Extractive Summarization using Deep Learning. *Research in Computing Science* 2018, 147, pp. 107-117.
18. Alami; Meknassi, M.; En-nahnah, N. Enhancing unsupervised neural networks based text summarization with word embedding and ensemble learning. *Expert Systems With Applications* 2019, 123, pp. 195-211.
19. Gulati, V.; Kumar, D.; Popescu, D.E.; Hemanth, J.D. Extractive Article Summarization Using Integrated TextRank and BM25+ Algorithm. *Electronics* 2023, 12, pp. 1-17.
20. Mojrian, M.; Mirroshande, S.A. A novel extractive multi-document text summarization system using quantum-inspired genetic algorithm: MTSQIGA. *Expert Systems With Applications* 2021, 171, pp. 1-20.
21. Dragomir-R, R.; Hongyan, J.; Małgorzata, S.; Daniel, T. Centroid-based summarization of multiple documents. *Information Processing & Management* 2004, 40, pp. 919-938.
22. Leonhard, H. Topic-based multi-document summarization with probabilistic latent semantic analysis. In: *Proceedings of the Proceedings of the International Conference RANLP-2009*, Borovets, Bulgaria, 2009, pp. 144-149.
23. Pengjie, R.; Zhumin, C.; Zhaochun, R.; Furu, W.; Liqiang, N.; Jun, M.; Maarten, D.R. Sentence relations for extractive summarization with deep neural networks. *ACM TOIS* 2018, 36, pp. 1-32.

24. Akülker, E. Extractive text summarization for Turkish using TF-IDF and pagerank algorithms. *Springer, cham.* 2022, 544, pp. 688-704.
25. Jia, R.; Cao, Y.; Tang, H.; Fang, F.; Cao, C.; Wang, S. Neural Extractive Summarization with Hierarchical Attentive Heterogeneous Graph Network. In: *Proceedings of the Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing*, online, November 16–20, 2020, pp. 3622–3631.
26. Rizvi, S.M.H.; Imran, R.; Mahmood, A. Text Classification using Graph Convolutional Networks: A Comprehensive Survey. *ACM Computing* 2025, 9, pp. 1-35.
27. Sijtsma, K.; Emons, W. Nonparametric statistical methods. *International encyclopedia of education* 2010, 7, pp. 347-353.

Citation: Shehu, A.; Karatu, M.; Roko, A.; Abdullahi, S. A graph convolutional network and recurrent neural network ensemble for extractive text summarisation in the Hausa language. *Journal of Engineering Science*. 2025, XXXII (1), pp. 32-46. [https://doi.org/10.52326/jss.utm.2025.8\(2\).03](https://doi.org/10.52326/jss.utm.2025.8(2).03).

Publisher's Note: JES stays neutral with regard to jurisdictional claims in published maps and institutional affiliations.



Copyright:© 2025 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

Submission of manuscripts:

jes@meridian.utm.md