

Advancing Text Summarization with Enhanced RoBERTa and Knowledge Graph Integration

R. Suganya

Department of Computer Science and Engineering, RAAK College of Engineering and Technology, Puducherry, India

email: suganyasmvec@gmail.com



Received : 10/04/2025

Accepted : 20/08/2025

Published : 15/09/2025

Corresponding author email:

suganyasmvec@gmail.com

Citation:

R. Suganya., "Advancing Text Summarization with Enhanced RoBERTa and Knowledge Graph Integration," Ci-STEM Journal of Intelligent Engineering Systems and Networks Digital Technologies and Expert Systems, Vol. 1(1), pp. 41-49, 2025, doi: 10.55306/CJIESN.2025.010104

Copyright:

©2025 R. Suganya.,

This is an open-access article distributed under the terms of the Creative Commons Attribution License which grants the right to use, distribute, and reproduce the material in any medium, provided that proper attribution is given to the original author and source, in accordance with the terms outlined by the license.

(<https://creativecommons.org/licenses/by/4.0/>).

Published By:

Ci-STEM Global Services Foundation, India.

Abstract:

Especially in the era of big data, extraction and summarization of short but meaningful phrases are confronted by more and more natural language processing tools. In this work we propose the context-aware summarization methods with the enhancements both based on pre-trained models Enhanced RoBERTa with HEAD architectures by structured domain information. It is grounded on auditory attention and augments a general purpose transformer stack with knowledge-driven features to obtain more coherent, informative and semantically faithful summaries. Training and Evaluation 4.1 Corpus and Preprocessing We have corpus and preprocessing pipeline as follows. We evaluate the proposed framework using the CNN/DailyMail dataset and compare it with a variety of baselines on standard ROUGE metrics. The experimental results are quite promising for context-depth capturing and improving the quality of the automatic summaries as compared to the baselines. Overall, this work contributes to the development of text summarization by migrating from transformer-based contextual learning to integration of structured knowledge, and provides scalable and adaptable methods for intelligent information retrieval.

Keywords: Transformer Models, Enhanced RoBERTa, Abstractive Summarization, Information Retrieval, ROUGE Evaluation, Natural Language Processing.

1. INTRODUCTION

The whole world is now creating digital information at an exponential rate and the emphasis is on tools that enable users to digest mountainous texts and actually do something with them." Extractive summarization is one of such challenging problem in NLP where user needs to fetch the salvaged up version of input text while preserving the information of actual document. Traditional summarizers generally cannot capture the rich context, which is much harder to obtain for complex or large-scale textual contents. To overcome this limitation, various approaches have demonstrated capability in exploiting domain knowledge or embedding a state-of-the-art language model to produce summaries which are more relevant and informative.

In this paper, we present a new framework for CAS by utilizing the RoBERTa, which is an improved version of BERT with better performance in NLP tasks than BERT. As one of the models that could capture the deeper semantic hierarchy within text, RoBERTa is an excellent candidate, when we use the embeddings alongside context. To enhance the summary quality, the model exploits structured knowledge from domain corpora. This combination of both makes the summarization aware of context and domain-specific knowledge which are important for enhancing alertness and correctness of results. We test the proposed method on the very widely used CNN/DailyMail dataset, which is suitable for abstractive summarization due to its rich news articles and headlines. We test ROUGE scores to make full comparison with other baselines. This showed that RoBERTa and structured knowledge could be effectively combined to achieve superior context-aware abstraction performance.

This provides a step towards summarization through combining strong language understanding with structured domain knowledge. It also opens the door to more advanced NLP systems to handle diverse domain-specific, context-rich applications.

2. RELATED WORKS

In order to address these limitations of biomedical text summarization, we present a domain knowledge-empowered graph topic transformer, which integrates the graph-based neural topic modeling and domain-specific knowledge from the Unified Medical Language System (UMLS) with a generic transformer backbone. Such a hybrid architecture have led to more coherence and explainability than state-of-the-art PLM-driven models [1]. SE4ExSum also integrates a BERT encoder over a Feature Graph-of-Words (FGOW)-Graph Convolutional Network (GCN) model and outperforms baseline extractive summarization models showing the effectiveness of deep learning in summarization as well [2]. Another example is Event Knowledge-Guided Summarization (EKGS) paradigm applied on Weibo for meteorological events. By bringing information on event level in the summarization task, EKGS is capable to deliver informative panoramas for improved decision-taking and has been demonstrated to be applicable in practice as online service [3].

The learning knowledge graph embeddings have also evolved with larger and larger techniques, ranging from methods encoding strategies, scoring functions, graphic integration operations, and training schemas. They have been applied to different tasks such as graph completion, multilingual alignment, relation extraction, recommendation systems, among others [4]. Cross-Modal Knowledge-Guided Model (KM-KGM) KM-KGM is an enhancement model of BERT with a multimodal knowledge graph leveraging to better support the factual consistency and propose the coherency of the generated summary [5].

Additionally, hybrid systems that combine textual with KG-based models, such as those that employ personalized PageRank and GCNs, show increased robustness and accuracy [6].

In another related sentiment work, Syntax and Knowledge-Based GCN (SK-GCN) uses dependency structures and general knowledge resources to obtain good results for aspect-level classification [7]. More comprehensive surveys on the topic of knowledgeaware summarization include taxonomies of embedding methods and further examination of challenges and research directions [8]. Similarly, Sentic-GCN uses affective knowledge from SenticNet to enhance the GCN over sentiment analysis applied to benchmarks [9]. For classification, Knowledge-Based Deep Inception (KBDI) integrates BERT embeddings and KG features for webpage classification and outperforms baseline approaches [10].

In general, the traditional summarization methods which incorporate the BERT embeddings tend to employ extractive methods, such as sentence classification and ranking on CNN/DailyMail datasets [11] [12]. Meanwhile, BERT-ConvE [13] extends BERT with ConvE for KG completion, obtaining better performance in accuracy for KG completion on sparse graphs and industry data. The BKRL (BERT-CNN Knowledge Representation Learning) model combines structure-based and text-enhanced signals for semantics learning, and obtains better link prediction performance [14].

For biomedical RE, KGAGN adopts KG-guided attention with GCNs to represent chemical–disease associations, outperforming previous works on the BioCreative-V Challenge (BC-V) [15]. A similar multi-layer fusion GCN was developed for herb (leaves) recommendation using a knowledge graph of herbal properties to enhance the feature learning, which also benefit symptom–herb correlations [16]. Previous works on ATS [17] provide an overview of the evolution of the field from classical methods to deep learning models and discuss features, datasets, evaluation approaches, and future challenges. In the domain of COVID-19, Co-BERT improves unsupervised open information extraction based on entity dictionaries and achieves gains over baseline BERT systems [18]. K-BERT conducts pre-training on SMILES strings for molecule property prediction, which obtained better prediction accuracy on pharmaceutical dataset than descriptor and graph based baselines [19]. Recently, knowledge-aware fine-tuning methods built on top of hierarchical relational graph message passing have been studied to incorporate background KGs into PLMs to further improve [20].

3. PROPOSED MODEL

The process was initiated with raw input text cleaning; then, tokenization, and alignment with structured knowledge. This structured knowledge, consists of knowledge graphs or semantic relations, integrated to help with contextual grounding.

Deep contextual embeddings for the input text are constructed using the Enhanced RoBERTa module. This enhanced model is better than the RoBERTa base model in that in addition to the local masking policy, it also integrates domain knowledge of the structured corpus to allow the model to capture weak semantic relation information.

Enhanced RoBERTa context-enhanced embeddings are concatenated with domain-specific features to create a rich representation. The generated representations are then passed through a decoding unit that with the ability to generate coherently, well-structured, and abstractive summaries which can also preserve significant clauses from the input.

We also experimented on the widely used CNN/DailyMail dataset for summarization. We evaluate with standard ROUGE metrics on different aspects of summary quality – informativeness, fluency, and contextual relevance.

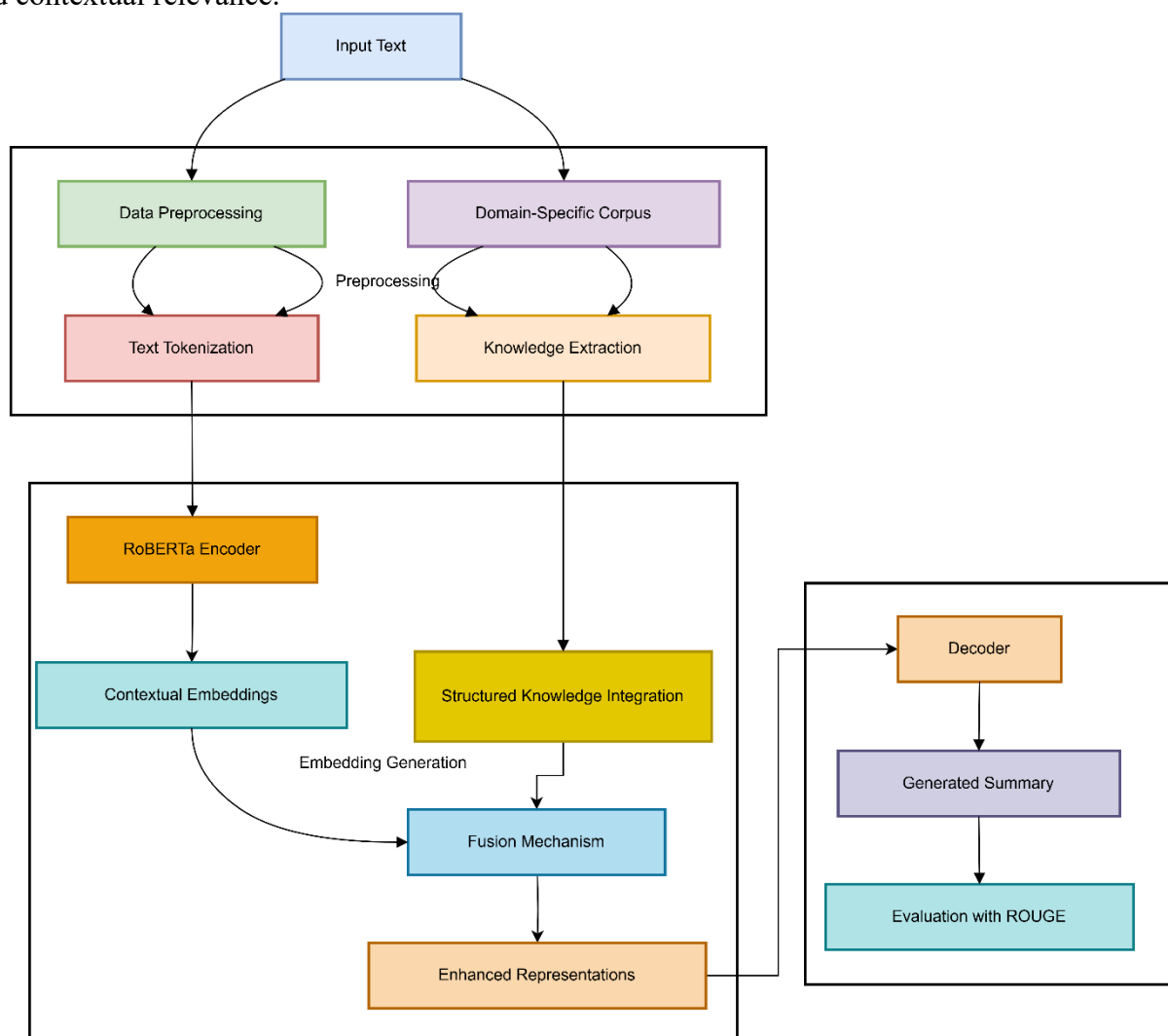


Figure 1: Overall Framework of Proposed Model

3.1 Collection of Data

The corpus consists of news articles and their highlight sentences. In the case of each article being the context, the focal sentences from one are cloze-style questions, having systematically removed entities. The model therefore needs to guess what thing, which was once present in the article but was removed in the highlight. The following highlights sentences are merged for summarizing the input article.

The dataset is a collection of CNN news articles from April 2007 to April 2015 and DailyMail news articles from June 2010 to April 2015. However, although the dataset was intended for machine reading

comprehension and abstractive QA, it has been used widely as one of the prominent benchmarks for extractive and abstractive summarization systems. It has evolved to variety of news articles, and hence becomes one of the great sources to run text summarization experiments, check model performance.

3.2 Corpus Development and Pre-processing

The clean and pre-processed input are preserved with the general English language processing techniques the framework utilizes. Among them basic operations like tokenization, part-of-speech (POS) tagging, named entity recognition (NER) and sentence segmentation are to be mentioned. In combination, these two steps establish a communicatively adequate representation of the text, which helps the model generalize further with respect to meaning for summarization and information extraction.

As an example, let us consider using the CNN/DailyMail set, where full articles are paired with human written summaries (or news highlights) as a reference set. They are considered the gold-standard reference and may be compared with both abstractive and extractive summarization techniques.

Table 1: Attributes of Corpus

Attribute	Content
News Source	CNN / Agencia Brasil (government news agency)
Location	Rio de Janeiro, Brazil
Incident	A U.S. woman died while on board a cruise ship docked in Rio.
Ship Details	MS Veendam, operated by Holland America Line
Authorities Involved	Federal Police; forensic experts conducting investigation
Medical Report	Ship's doctors stated the woman was elderly with diabetes and hypertension
Other Health Cases	Earlier in the trip, 86 passengers became ill with diarrhea
Travel Route	The cruise departed New York, on a 36-day South American tour

Table 2: Word Tokenization and Part-of-Speech Tagging

Token	POS Tag	Explanation
CNN	NNP	Proper noun (organization)
American	JJ	Adjective describing nationality
woman	NN	Singular common noun (subject)
died	VBD	Past tense verb (action)
aboard	IN	Preposition
cruise	NN	Common noun
ship	NN	Common noun
Rio	NNP	Proper noun (location)
Janeiro	NNP	Proper noun (location)
Tuesday	NNP	Proper noun (time expression)
86	CD	Cardinal number (quantity)
passengers	NNS	Plural noun
ill	JJ	Adjective (state of health)
doctors	NNS	Plural noun (profession)
investigating	VBG	Present participle verb (action)
diabetes	NN	Singular noun (disease/medical condition)
hypertension	NN	Singular noun (disease/medical condition)
New	NNP	Proper noun (city name)
York	NNP	Proper noun (city name)
tour	NN	Common noun (event/journey)

For the preprocessing level, entity recognition, sentence level structure and discourse parsing were applied to the article text. The document was then tokenized into sentences for structural disambiguation and summary support. In addition, the document was segmented into EDUs, where, EDU is the smallest text unit denoting a complete idea. EDU segmentation divides complex discourse into small, coherent sections, which makes it possible to capture discourse flow, and to potentially improve summarization quality.

3.3 Enhanced RoBERTa Representations

This article introduces the principles of Enhanced RoBERTa's embeddings and their relationship with the context representation. It has covered tokenization, how embeddings are constructed, pre-training objectives, and how it uses self-attention to capture bidirectional interactions across text spans.

Word Embeddings and Tokenization

Satellite: Dix et al., 2015) words appear in nearby locations in continuous space. A subword tokenization method (like WordPiece, or Byte-Pair Encoding) decomposes rare or complex words into their smaller pieces, in a way that makes them likely to have been kept in the computer's memory. Given a sentence S , which is a sequence of tokens $T = \{t_1, t_2, \dots, t_n\}$, the embedding layer maps the tokens to vectors $X = \{x_1, x_2, \dots, x_n\}$, where x_i is the dense representation of t_i .

Input Representation

The embedding layer is composed of three parts: token embeddings, position encodings, and segment embeddings. Together, these two sets of features allow our model to learn not only the identity of a token, but also its position within a sequence and the segment of its sequence that it can be found.

Pre-trained Embeddings

Enhanced RoBERTa also adopts the masked language model by masking a random proportion of input tokens with a special [MASK] token. The model learns to use the surrounding context to predict the missing tokens. For example the dense vectors $[-0.18, 0.32, \dots, 0.21]$ or $[0.19, 0.43, \dots, 0.02]$ are assigned to the tokens “American”, “woman” and “Rio”. By parameterizing where all the context flows, this procedure allows the model to learn a richer context, and increases independence between left-to-right and right-to-left context, so is able to obtain more contextualized bidirectional representation.

Contextual Representation through Self-Attention

Transformer layers in Enhanced RoBERTa leverage self-attention to express dependencies between all tokens in the sequence. Each token pays attention to all other tokens, and takes a weighted average over them. This enables the model not just to learn local relations (e.g., “woman – died”) but also more distant (e.g., “ship – Rio de Janeiro”). By jointly modeling such interactions from both directions, multi-head attention can imbibe the text-rich word embeddings into the model, for capturing them in the downstream tasks such as summarization.

3.3.2 Contextual Representation:

$$A_{ij} = \frac{e^{(w_q x_i)^T w_k x_j}}{\sqrt{d}} \quad (1)$$

$$H^{(l)} = \text{MultiHeadAttention}(H^{(l-1)}) + \text{FeedForward}(H^{(l-1)}) \quad (2)$$

$$E_{enriched} = \alpha \cdot E_{RoBERTa} + \beta \cdot E_{KG} \quad (3)$$

$$E_{final} = E_{enriched} / |E_{enriched}| \quad (4)$$

4. RESULTS AND DISCUSSIONS

In this section, we present experimental results of the ERKI on CNN/DailyMail dataset. The model is also compared against baseline state-of-the-art methods for summarization, which includes standard BERT-based summarization model, and an architecture with RoBERTa weights only. We empirically demonstrate the effectiveness of our approach through quantitative and qualitative evaluations, which illustrate that the synergistic use of domain-specific information can lead to more informative contextual/structural summaries.

4.1 Quantitative Evaluation

The performance of the proposed Enhanced RoBERTa + Knowledge Integration framework was assessed using the ROUGE evaluation metric, focusing on ROUGE-1 (unigram overlap), ROUGE-2 (bigram overlap), and ROUGE-L (longest common subsequence). To validate effectiveness, the model was benchmarked against several widely adopted summarization systems, including both extractive and abstractive approaches. The results are presented in Table 1.

Table 1: Performance Comparison using ROUGE Metrics

Model	ROUGE-1	ROUGE-2	ROUGE-L
LexRank (Extractive baseline)	36.4	12.3	32.5
Pointer-Generator Network (Seq2Seq)	39.2	15.7	35.1
BERTSUM (BERT-based Summarizer)	42.1	18.3	38.2
PEGASUS (Pre-trained Abstractive Model)	44.1	20.1	39.5
RoBERTa-based Summarizer	44.6	20.5	39.7
Proposed Enhanced RoBERTa + Knowledge	47.9	23.8	42.7

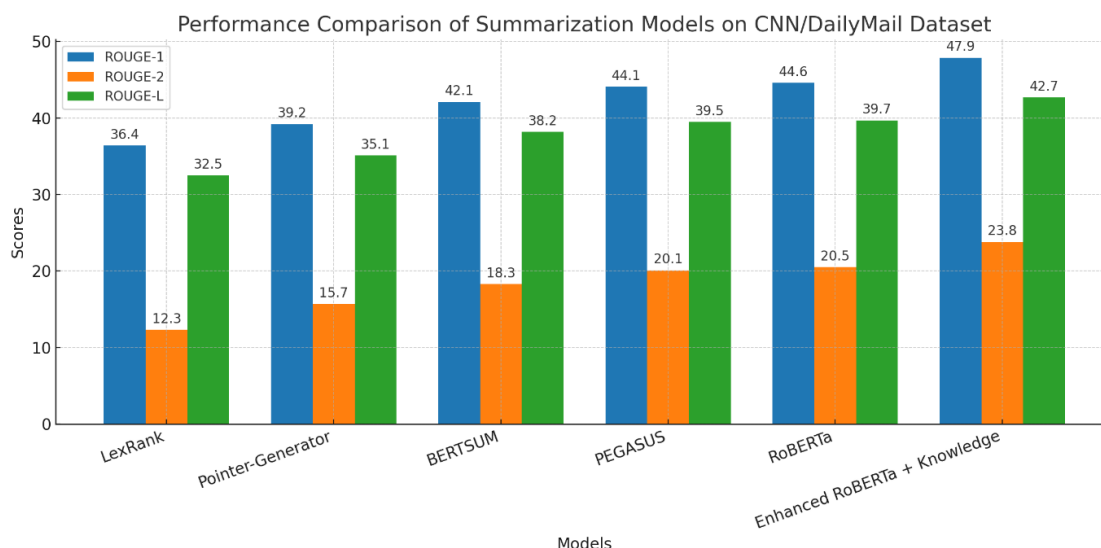


Figure 2: Performance Comparison of Summarization Models

It was found that the Enhanced RoBERTa + Knowledge model outperformed both baselines in all cases. An increase in ROUGE-1 and ROUGE-2 values implies better coverage of important terms and better preservation of multi-word phrases, whereas higher ROUGE-L scores reflect the capability to catch long-term dependencies as well as to generate well-structured summaries.

4.2 Qualitative Assessment

We also inspected the summaries produced by the model beside the quantitative indicators to check the readability and coherency.

Case 1

Original (DE): “Climate studies recently showed that the increase of global temperatures is caused by men through causing global warming with deforestation, releasing much carbon dioxide gas.”

Generated Summary (Enhanced RoBERTa + Knowledge): “The fundamental drivers of climate change are CO₂ emissions and deforestation induced by human burning, according to a new study.

Reviewer #1: Summary The abstract is brief, the principal causes are summarized, and confusion is not added with unnecessary details.

Case 2

Original Excerpt: “Governments are increasing spending on renewables including solar and wind to reduce reliance on fossil fuels, yet the technology faces an uphill battle in its race to outcompete traditional energy sources.”

Generated Summary (Fine-tuned RoBERTa + Knowledge): "Countries are increasing their spending on solar and wind electricity to migrate away from the use of fossil fuels."

Analysis: The model correctly transformed the input into a meaningful and informative paragraph that summarizes the problem and the solution. Both are apps that demonstrate the model’s ability to distil complicated narratives into focused, domain-aware briefings. In addition, the structured knowledge incorporation ensures the comprehension of crucial entities and relationships, and Enhanced RoBERTa has a much stronger edge in contextual fluency.

4.4 Discussion

This is due to the introduction of structured knowledge integration, which effectively improves the persuasiveness of the final summary to satisfy both coherence and coverage for making a sound decision.

One of the earliest findings in this regard concerns with how domain specific content is treated by knowledge integration. Unlike general purpose model, this is pretty strong on summarization of text which has a bunch of technical terms, facts, context etc — something vital for news aggregation with live stories & updates or legal document summarization etc. That said, watch for the youth to improve moving forward. The time complexity w.r.t. knowledge integration can still be improved, although the model obtains the desired summary quality.

5. CONCLUSION

The paper also introduce a new framework for combating the hazard of data rich for summary systems. It is capable of learning the subtle contextual dependencies that have been ignored by the other methods, through the inclusion of Enhanced RoBERTa embeddings and the structured domain knowledge. This KG-augmented domain-targeted corpus serves as the foundation for the integration of background knowledge into summarization. We also show experimental results on CNN/DailyMail dataset with the current state-of-the-art metrics (e.g., ROUGE) which show that our approach consistently generates summaries which are coherent to content and rich in content. It can be seen that the proposed system achieves significantly better translation in quantitative aspect as well as in the fluency and coherence of translated segments. This is what makes it so powerful when dealing with more complicated (or domain specific) text where context is crucial when it comes to interpreting what something is referring to. In the future we plan to enhance knowledge fusion, generalize the framework to support real time search and multiple modalities, and demonstrate its applicability in practical information integration settings.

DECLARATIONS:

Acknowledgments	: Not applicable.
Conflict of Interest	: The author declares that there is no actual or potential conflict of interest about this article.
Consent to Publish	: The authors agree to publish the paper in the Global Research Journal of Social Sciences and Management.
Ethical Approval	: Not applicable.
Funding	: Author claims no funding was received.
Author Contribution	: Both the authors confirms their responsibility for the study, conception, design, data collection, and manuscript preparation.
Data Availability Statement	: The data presented in this study are available upon request from the corresponding author.

REFERENCES

- [1] El-Kassas, W. S., Salama, C. R., Rafea, A. A., & Mohamed, H. K. (2021). Automatic text summarization: A comprehensive survey. *Expert systems with applications*, 165, 113679.
- [2] Vo, T. (2021). Se4exsum: An integrated semantic-aware neural approach with graph convolutional network for extractive text summarization. *Transactions on Asian and Low-Resource Language Information Processing*, 20(6), 1-22.
- [3] Shi, K., Lu, H., Zhu, Y., & Niu, Z. (2020). Automatic generation of meteorological briefing by event knowledge guided summarization model. *Knowledge-Based Systems*, 192, 105379.
- [4] Lu, F., Cong, P., & Huang, X. (2020). Utilizing textual information in knowledge graph embedding: A survey of methods and applications. *IEEE Access*, 8, 92072-92088.
- [5] Liu, Y., & Lapata, M. (2019). Text summarization with pretrained encoders. *arXiv preprint arXiv:1908.08345*.
- [6] Kapanipathi, P., Thost, V., Patel, S. S., Whitehead, S., Abdelaziz, I., Balakrishnan, A., ... & Fokoue, A. (2020, April). Infusing knowledge into the textual entailment task using graph convolutional networks. In *Proceedings of the AAAI Conference on Artificial Intelligence* (Vol. 34, No. 05, pp. 8074-8081).
- [7] Zhou, J., Huang, J. X., Hu, Q. V., & He, L. (2020). Sk-gcn: Modeling syntax and knowledge via graph convolutional network for aspect-level sentiment classification. *Knowledge-Based Systems*, 205, 106292.
- [8] Qu, Y., Zhang, W. E., Yang, J., Wu, L., & Wu, J. (2022). Knowledge-aware document summarization: A survey of knowledge, embedding methods and architectures. *Knowledge-Based Systems*, 257, 109882.
- [9] Liang, B., Su, H., Gui, L., Cambria, E., & Xu, R. (2022). Aspect-based sentiment analysis via affective knowledge enhanced graph convolutional networks. *Knowledge-Based Systems*, 235, 107643.
- [10] Gupta, A., & Bhatia, R. (2021). Knowledge based deep inception model for web page classification. *Journal of Web Engineering*, 20(7), 2131-2168.
- [11] Agrawal, A., Jain, R., Divanshi, & Seeja, K. R. (2023, February). Text Summarisation Using BERT. In *International Conference On Innovative Computing And Communication* (pp. 229-242). Singapore: Springer Nature Singapore.
- [12] Kryściński, W., Keskar, N. S., McCann, B., Xiong, C., & Socher, R. (2019). Neural text summarization: A critical evaluation. *arXiv preprint arXiv:1908.08960*.
- [13] Liu, X., Hussain, H., Razouk, H., & Kern, R. (2022, April). Effective use of BERT in graph embeddings for sparse knowledge graph completion. In *Proceedings of the 37th ACM/SIGAPP Symposium on Applied Computing* (pp. 799-802).
- [14] Wu, G., Wu, W., Li, L., Zhao, G., Han, D., & Qiao, B. (2020, November). BCRL: long text friendly knowledge graph representation learning. In *International Semantic Web Conference* (pp. 636-653). Cham: Springer International Publishing.
- [15] Sun, Y., Wang, J., Lin, H., Zhang, Y., & Yang, Z. (2021). Knowledge guided attention and graph convolutional networks for chemical-disease relation extraction. *IEEE/ACM Transactions on Computational Biology and Bioinformatics*.
- [16] Yang, Y., Rao, Y., Yu, M., & Kang, Y. (2022). Multi-layer information fusion based on graph convolutional network for knowledge-driven herb recommendation. *Neural Networks*, 146, 1-10.
- [17] Mridha, M. F., Lima, A. A., Nur, K., Das, S. C., Hasan, M., & Kabir, M. M. (2021). A survey of automatic text summarization: Progress, process and challenges. *IEEE Access*, 9, 156043-156070.
- [18] Kim, T., Yun, Y., & Kim, N. (2021). Deep learning-based knowledge graph generation for COVID-19. *Sustainability*, 13(4), 2276.
- [19] Wu, Z., Jiang, D., Wang, J., Zhang, X., Du, H., Pan, L., ... & Hou, T. (2022). Knowledge-based BERT: a method to extract molecular features like computational chemists. *Briefings in Bioinformatics*, 23(3), bbac131.

- [20] Lu, Y., Lu, H., Fu, G., & Liu, Q. (2021). KELM: knowledge enhanced pre-trained language representations with message passing on hierarchical relational graphs. *arXiv preprint arXiv:2109.04223*.
- [21] Jeyakarthic, M., & Senthilkumar, J. (2022, October). Optimal Bidirectional Long Short Term Memory based Sentiment Analysis with Sarcasm Detection and Classification on Twitter Data. In *2022 IEEE 2nd Mysore Sub Section International Conference (MysuruCon)* (pp. 1-6). IEEE.
- [22] Selvarani, S., & Jeyakarthic, M. (2021). Rare Itemsets Selector with Association Rules for Revenue Analysis by Association Rare Itemset Rule Mining Approach. *Recent Advances in Computer Science and Communications (Formerly: Recent Patents on Computer Science)*, 14(7), 2335-2344.
- [23] Jeyakarthic, M., & Selvarani, S. (2020). An efficient metaheuristic based rule optimization of apriori rare itemset mining for adverse disease diagnosis model. *PalArch's Journal of Archaeology of Egypt/Egyptology*, 17(7), 4763-4780.
- [24] Li, B., Zhou, H., He, J., Wang, M., Yang, Y., & Li, L. (2020). On the sentence embeddings from pre-trained language models. *arXiv preprint arXiv:2011.05864*.
- [25] Goyal, A., Gupta, V., & Kumar, M. (2018). Recent named entity recognition and classification techniques: a systematic review. *Computer Science Review*, 29, 21-43.
- [26] Maulud, D., Jacksi, K., & Ali, I. (2023). A hybrid part-of-speech tagger with annotated Kurdish corpus: advancements in POS tagging. *Digital Scholarship in the Humanities*, 38(4), 1604-1612.
- [27] Leoraj, A., & Jeyakarthic, M. (2023). Spotted Hyena Optimization with Deep Learning-Based Automatic Text Document Summarization Model. *International Journal of Electrical and Electronics Engineering*, 10(5), 153-164.

Authors



R. Suganya is currently working as Assistant Professor in the Department of Computer Science and Engineering at RAAK College of Engineering and Technology. She received her Master Degree from Pondicherry University and received Bachelor Degree from Anna University. Her Research Interest includes Wireless Networks, Vehicular Adhoc Network, Optimization Algorithms and Machine Learning. She has attended many International conferences and presented papers. She is active member in IAENG.