

Comprehensive Analysis of Text Summarization Techniques for Legal Documents

Ryakala Deepika¹, Surajit Das², Niladri Sekhar Dey³, Jeshwanth Panuganti⁴ and Mohammed Raashed Hussain⁵

¹⁻⁴ B V Raju Institute of Technology, Dept. of AI&DS, Narsapur, India

Email: deepika.r@bvrit.ac.in, surajit.das@bvrit.ac.in, niladri.dey@bvrit.ac.in

⁵ B V Raju Institute of Technology, Dept. of CSE (AI&ML) Narsapur, India

Email: raashedhussain315@gmail.com

Abstract— The surge of digital legal documents has significantly expanded their usage. This has resulted in the sheer number of papers being used by various members of the judiciary, as well as by advocates and judicial officers. It can be incredibly challenging to keep up with all of them. Over four crore cases are still pending in Indian courts, and manually reviewing them can be a tedious and time-consuming process. As machine learning has advanced, various text summarization models have been created to help legal professionals manage their documents. Due to the lack of publicly accessible datasets, it is difficult to fine-tune domain-independent models for Indian legal systems. The methodology proposed in this paper seeks to improve the overall performance of these models, and it also explores Indian legal documents' summarization techniques. In addition, this research also provides a study of the several summarization methods in-depth that have been on Indian legal documents, including PEGASUS, Bidirectional Auto-Regressive transformers (BART), TextRank, and Bidirectional Encoder Representations from Transformers (BERT). Through the process of extractive and abstract summarization, BART and PEGASUS will be able to gain a deeper understanding of the text normalization process. The outcomes of the text normalization process are evaluated by experts using the ROUGE metrics and multiple parameters. It shows that the proposed approach can work well in legal texts that have domain-independent frameworks.

Index Terms— Text-Rank, BERT, BART PEGASUS.

I. INTRODUCTION

Over the past few years, there has been a huge increase in the number of text resources found on various websites, like social network sites, Wikipedia, Legal Text Datasets, and content farm sites. Because of the growth of information, the process of properly utilizing these resources has become a major challenge [2]. A well-organized summary of a lengthy ruling may be just as informative and useful as reading the whole documents itself [3]. The process of automatic text summarization (ATS) involves selecting the most crucial ideas in a text to help the reader comprehend the documents. In India, a judgment is a legal decision that a court makes regarding a case. Due to the rapid emergence and evolution of new technologies, the documents related to this process are now stored digitally, making them available in a short time. This has made it easier for the law enforcers to keep track of the judgements [4]. Typically, a legal practitioner would begin preparing his statement about an issue by researching similar court documentation or rulings to utilise as references and as explanations presenting the appeal. These court documents are so wordy that sometimes include repetitive information [5, 6]. In the US alone, courts and the

laws create vast amounts of text each day. Caseworkers and judges handle millions of cases every year, and these file cases with thick legal content that can sometimes go several hundred pages long [7]. Therefore, that is useful to select sections that might also contribute towards the synopsis of the document. Thus, an overview might preserve the original document's key ideas while minimizing the vocabulary size, processing time, and space needs for any future analysis. Some method of automating or streamlining the review process might aid legal professionals in managing their workload. It is also beneficial for novices and average individuals to comprehend a judgment. Manually creating case summaries is a time-consuming operation. As a reason, having the material condensed into a concise form is both more practical for human readers and more advantageous to machine reading systems [8]. Accelerating advancements in areas including Artificial intelligence (AI) [9,10], data analysis and Deep Learning (DL) [11,12] made the automatic summarization of texts feasible [1]. DL techniques have recently made considerable strides in numerous Natural language processing (NLP) applications [13, 14]. In particular, Text Generation (TG) activities are among the hardest NLP jobs because they need automated text interpretation and precise semantic and lexical analysis. MT, IC, and ATS are examples [15, 16]. The two broad categories of summarization that are presented in automatic tools are abstractive and extractive summarization. These are respectively focused on the synthesis and examination of diverse textual materials [8]. The initial article subject's rank is assigned to every word in the text during the extractive summarization process. Top-scoring statements then summarize the document [17]. Text-Rank is a widely used method for extractive summarization. The goal of this study is to analyse the various strategies that can be used to improve the efficiency of text summarization in the context of legal texts. The scope and related work were discussed in section II, and the detailed technology usages in the text summarized was discuss in section III. Section IV represents the experience setup and obtained results and discussion were presented in section V. Finally, the concluding remarks with future recommendation were discuss in section VI.

II. RELATED WORKS

The usage of several domain-specific acronyms makes legal case judgements often long and convoluted. A lot of research is conducted on the legal text summarizing in different countries, such as the UK, Canada, and Denmark. For legal text summarization, the majority of research attempted to build supervised or semi-supervised DL techniques [18]. It was introduced by K. U. Manjar et al. [19] in 2020. In addition, various conventional extraction techniques, such as textual summaries, based on the terms' frequency analysis, were presented [20]. It pioneers the use of the graph model in the realm of autonomous summarizing. Methods for the graph-oriented summarization use phrases or paragraphs to make the "modules" inside a network, which are then ranked by how important they are and how much they are alike. Different graph-based techniques are additionally available for extracting summary [21], which often do well on textual summarization challenges. Great improvements in the domain of DL have led for development of extractive summarization algorithms that use DL techniques [22, 23]. Abstractive summarization has been studied using RNN-based DL, reinforcement learning, and pre-training models [24, 25]. Though modern NLP is propelled by models based on transformers for tackling sequence to sequence modelling issues, it appears there is a dearth of study when the topic at hand is the abstractive summary of legal discourse. A machine learning-based prediction model known as eLegPredict was developed by Sharma et al [26]. It was able to predict the cases decided by India's Supreme Court. The model was trained using X Gradient Boost classifier. The model was able to achieve an accuracy of 76 over 3072 judgements. When a new case is added to a directory, the prototype automatically goes through the details and provides a prediction. According to the study conducted by Pillai et al. [27], in order to predict a legal case, using convolutional neural network (CNN) [28,29] and NLP, they proposed the Bag of Words algorithm, which is a mechanism that allows users to select the words from the legal text. The CNN algorithm was then used to classify the cases into IPC section, and the outcome was predicted to be either a non-bailable or bailable case. The researchers were able to achieve an accuracy of 85% in the prediction of the outcome based on the IPC. They noted that the model performed better when the case only had one charge. The accuracy of their prediction decreased as the number of charges rose. According to Shaikh et al [30], the model was developed to predict the outcome of a case involving homicide. The researchers used different machine learning techniques to analyze and interpret the data collected from the Delhi district court. Some of the classification algorithms that were used in the study included the regression and classification tools known as CART, LR, SVM, and NB. The results of the experiment revealed that the classification system performed well in terms of both accuracy and F1 score. It was able to achieve an accuracy range of 85% to 92% But it did not perform well when it came to cases with more than one accused individual. El-Kassas et al. [31] developed a system in which "EdgeSumm" was assessed using the most popular automated assessment programme ROUGE and the famous short text summing datasets DUC2001 and DUC2002. This

evaluation uses ROUGE measurements, with EdgeSumm receiving the top ROUGE ratings on DUC2001. In his suggested framework, generic single-document summary has been implemented, and the assessment findings are quite encouraging. A. Joshi et al. [32] provide an overview of the various parameters that are involved in the creation of a comprehensive summary of a legal document. S. Polysley et al. [7] then developed a technology that automatically produces a summary based on the word frequency. Gupta et al. [33] then suggest that a customized summary can be extracted from online sources. The methods for legal text summarizing have been briefly mentioned by Nasar et al. [34], who have instead focused on discussing text summary in general. Our paper offers the full-fledged dissection for text summarization models for Legal Text Datasets by making use of preprepared language models such as BART, TextRank, BERT and PEGASUS models.

III. TECHNIQUES USED

A. Pegasus

It is a method similar to an extractive summary in that it extracts or masked important sentences from an input text and creates them as a single output sequence from the remaining phrases [25]. The figure 1 shows the architecture of the PEGASUS transformer-decoder base used for implementing various pre training objectives of the program. In this example, three sentences are used as targets for the pre-training objectives. The first sentence is masked using MASK1, while the other two are masked using MASK2. The purpose of this pre training is to improve the program's ability to understand and summarize text, particularly legal documents. The PEGASUS model is one of the summarization methods studied in the paper.

B. Text Rank

In this technique, a weighted network is generated using text units as nodes and a similarity metric based on word overlap to define edges, Similar to PageRank, this method determines the ranking of text units [35].

C. BERT

It's a model that's already been trained and allows us to do NLP tasks with ease. It is created to address the shortcomings of LSTM and RNN [36].

Disadvantages of LSTM

- LSTM are really susceptible to overfitting.
- It takes additional resources.
- However, LSTM cannot totally eliminate gradient issues

Disadvantages of LSTM

- It cannot properly manage lengthy sequences.
- High chance of vanishing gradient and ballooning gradient issues.
- Its training requires massive quantities of time and is challenging.

Advantages of BERT

- This pre-trained model needs no further training.
- The summary is based on key phrases from the larger text.
- It contains transformer-layered power encoder and decoder

D. BART

BART [37–39] is a transformer-oriented sequence model, that uses a denoising goal to learn appropriate approximations of input sentence for a variety of datasets. The model creates tokens from input sentence and use for machine translation and text summarization. Bart serves as both an encoder and a decoder. The encoder is responsible for extracting meaningful data from the provided text, while the decoder is used to determine the probability of the following word, so that the text may be rendered in a human-readable manner [40]. The BART algorithm architecture as shown in Figure 2. The model requires a data source, which is the input text that needs to be summarized. The work used the BART text summarization techniques for legal documents, including the use of machine learning models.

IV. EXPERIMENTAL SETUPS.

In this part, we go into depth about how we collected the dataset and how we evaluated its accuracy.

A. Dataset Preparation

The dataset contains case reports from the Federal Court of Australia, which were collected from 2006 to 2009. These reports were extracted from the AustLII free online legal database. This provides researchers with valuable insight into citation analysis and automatic summarization procedures. The goal of summarization experiments is

to explore the relationships between various elements. This dataset was collected from the public database repository of Kaggle [41].

B. Performance measures used

ROUGE [42] is a collection of performance criteria being used to assess different summarized texts and various automation transcriptions. It evaluates a computer-generated synopsis to a collection of standard summaries. ROUGE-N [37,42] It serves in assessing template matching, word embedding, and higher cognitive overlapping based just on n-gram overlaps. ROUGE-L [42] It determines the outputs of our algorithm in comparison to the benchmark and measures the frequent item sets subsequence of those outputs. Usually, summary is evaluated using ROUGE ratings or expert evaluation. We have determined the ROUGE-3, ROUGE-L, ROUGE-2, and ROUGE-1 scores. On the data, Recall, Precision, and F-measures, are calculated, and the superior of two strategies is advised for the IT LAW legal cases dataset.

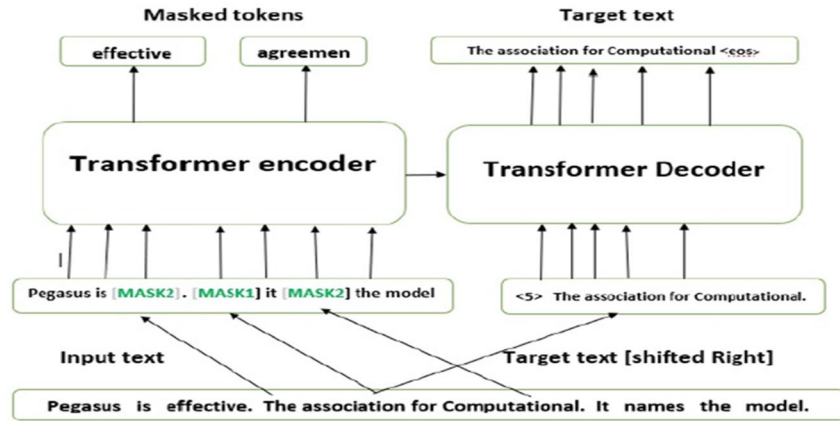


Fig. 1: In this example, three sentences are used as targets. The first sentence is masked using MASK1 while the other two are masked using MASK2 [25].

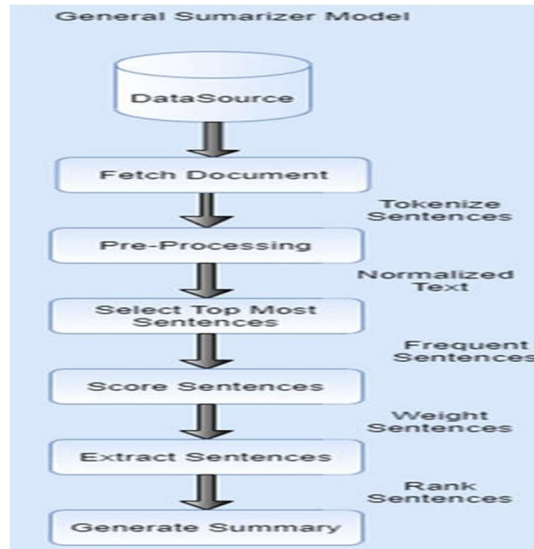


Fig. 2: Proposed BART model flowchart for text summarization

RECALL [37, 42, 43] It displays the proportion of human description covered by the conclusion reached. The formula for word overlap is as follows:

$$\text{Recall} = \frac{\text{The number of words that overlap}}{\text{Overall words in the citation abstract}} \quad (1)$$

Precision [37, 42] It evaluates the proportion of the system summary that is pertinent to the job. Formula for measuring accuracy in terms of overlapped of words:

$$\text{Precision} = \frac{\text{The number of words that overlap}}{\text{Overall words in the system abstract}} \quad (2)$$

F1-Score [37, 42, 44] This is the harmonic mean of recall and precision. The optimal value of f-measure is 1 while the worst value is 0.

$$\text{F1 - Score} = 2 * \frac{\text{Precision} * \text{Recall}}{\text{Precision} + \text{recall}} \quad (3)$$

V. RESULT AND DISCUSSION

The following section summarizes the major findings from the evaluation phase of the collected datasets.

TABLE I: RECALL SUMMARIZATION EVALUATION RESULTS

Model	Model	Rouge-2	Rouge-L
PEGASUS	0.357	0.313	0.357
Text-Rank	0.507	0.429	0.508
BART	0.534	0.462	0.536

TABLE II: PRECISION SUMMARIZATION EVALUATION RESULTS

Model	Model	Rouge-2	Rouge-L
PEGASUS	0.939	0.887	0.939
Text-Rank	1.0	1.0	1.0
BART	0.957	0.925	0.957

Python language utilised throughout the implementation of each and every method. The experiments were performed on three different models: BART, Text Rank, and PEGASUS. When it comes to summarization, BART performs better than its three counterparts. One of the main factors that contributed to this is the model's multi-head attention, which helps it learn a sentence better. This is because it allows it to pay attention to various parts of a sentence. The experiment Learning rate and Loss model performance Graph of BART is presented in Figure 3.

TABLE III: F1-SCORE SUMMARIZATION EVALUATION RESULTS

Model	Model	Rouge-2	Rouge-L
PEGASUS	0.517	0.462	0.517
Text-Rank	0.673	0.601	0.673
BART	0.678	0.616	0.687

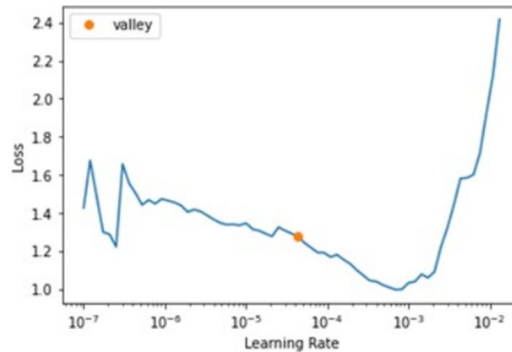


Fig. 3: Curves depicting the Learning rate vs Loss for the suggested model

Text	Target	Prediction
It is not a necessary consequence of an order appointing a receiver that the receiver should deal with or liquidate the assets in question. The interlocutory and protective character of orders made under s 1323 must be borne in mind when defining the powers of the receiver. The appointment of a receiver has rightly been described as 'an extraordinary step' ASIC v Burke [2000] NSWSC 694 at [8] (Austin J). However depending upon the nature of the powers conferred on the receiver it may be less drastic than a freezing order which can only be varied by order of the Court. The interlocutory history of this case has already demonstrated that circumstances not contemplated when the original interim freezing orders were made have required their variation from time to time. I accept, with respect, the observation made by Austin J in Burke at [8]: 'Without wishing to lay down any general rules, it appears to me that the extraordinary step of appointing a receiver may be justified, even though Mareva Orders are in place, in a case where there is real doubt about the existence and location of assets such as investments, and about the number and identity of claimants and the nature of their claims, and additionally the defendants are engaged in business activities which entail that any Mareva Orders must	ASIC v Burke [2000] NSWSC 694 at [8] (Austin J) ; [14] ACSC 28 ; [15] 14 ACSR 580 ; [16] 14 SCR 580; [17] 13SR 580; (14) 14 SCRA 355; (15) 13SR 600, Appellant S395/2002 v Minister for Immigration and Multicultural Affairs [2003] HCA 71 ; (2003) 216 CLR 473 (Appellants S396/2002 ; (2013) 14 ACSR 580 ; (13 ACSR 550 ; (2014) 14 CLR 580	

Fig. 4: Experiment Results (Bart)

BART, Text Rank, and PEGASUS were utilized to process the data collected from Australian court cases. The ROUGE system then performed a comprehensive analysis of the recall data and determined the recall score for each category. The corresponding scores were then determined using the system's optimized strategy. The evaluation was carried out by comparing the summarized data from the model with the results from the experts who were responsible for the analysis.

The table 1 shows the recall summarization evaluation results for three models: PEGASUS, Text-Rank, and BART. The evaluation metrics used are Rouge-and Rouge-L. PEGASUS has the lowest recall score for both Rouge-2 and Rouge-L, with scores of 0.357 and 0.313, respectively. Text-Rank has a higher recall score than PEGASUS, with scores of 0.507 for Rouge-2 and 0.429 for RougeL. BART has the highest recall score among the three models, with scores of 0.534 for Rouge-2 and 0.536 for Rouge-L. So, BART is the best model among the three for recall summarization, as it has the highest scores for both Rouge-2 and Rouge-L. PEGASUS has the lowest recall scores, indicating that it may not be the best choice for recall summarization. Text-Rank has moderate recall scores, but it is still outperformed by BART. Table 2 shows the precision summarization evaluation results of three models: PEGASUS, Text-Rank, and BART. The evaluation metrics used are Rouge-2 and Rouge-L. PEGASUS has a Rouge-2 score of 0.939 and a Rouge-L score of 0.939. Text-Rank has a perfect score of 1.0 for both Rouge-2 and Rouge-L. BART has a Rouge-2 score of 0.957 and a Rouge-L score of 0.957. Text-Rank has the highest precision score among the three models. PEGASUS and BART have similar precision scores, with BART having a slightly higher Rouge-2 score. All three models have high precision scores, indicating that they are effective in summarizing legal documents. Table 3 shows the F1-Score summarization evaluation results for three models: PEGASUS, Text Rank, and BART. The table includes the Rouge-2 and Rouge-L scores for each model. BART has the highest F1-Score of 0.678, followed closely by Text-Rank with a score of 0.673. PEGASUS has the lowest F1-Score of 0.517. The Rouge-L scores are higher than the Rouge-2 scores for all three models. BART and Text Rank are better models for summarization compared to PEGASUS. Rouge-L is a better metric for evaluating summarization performance than Rouge-2. BART has the highest F1-Score among the three models, indicating that it is the most effective model for summarization.

The paper's analysis suggests that the additional comparative test's score is inferior to that of the BART framework. BART's results were the most outstanding in the tests. The results of this BART experiment can be seen in Figure 4. The Figure has three columns: Text, Target, and Prediction. The Text column contains a long sentence from a legal document. The Target column contains the summary of the sentence in the form of a case citation. The table shows the performance of the BART model on a specific sentence from a legal document. The Target column provides a reference to a relevant case, which can help legal professionals quickly understand the context of the sentence.

BART is the most efficient model when it comes to addressing the task of summary generation in terms of both performance and efficiency. Based on the results, it can be regarded as an effective and efficient method for tackling the abstract summarization problem.

VI. CONCLUSIONS

The text summarization research field is a fascinating area of study that involves analysing the various applications of web scrapping technique. This process is mainly used to extract online cases. It is additionally utilized to analyse the performance of different algorithms. The goal of the research is to help the various lawyers gain a better understanding of the information presented in their large documents. Through the analysis of the results, the BART is the best algorithms for summarization were identified. Our method has several benefits, but it also has some drawbacks which we want to address in our next works. Our model uses the summary output length as a parameter, thus even if a phrase is incomplete, it will chop off any excess tokens created outside the token length. To prevent incomplete sentences, build a longer output and leave off the final sentence.

REFERENCES

- [1] Satyajit Ghosh, Mousumi Dutta, and Tanaya Das. Indian legal text summarization: A text normalisation-based approach. arXiv preprint arXiv:2206.06238, 2022.
- [2] Ming-Hsiang Su, Chung-Hsien Wu, and Hao-Tse Cheng. A two-stage transformer based approach for variable-length abstractive summarization. *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, 28:2061–2072, 2020.
- [3] Ayesha Sarwar, Seemab Latif, Rabia Irfan, Adnan Ul-Hasan, and Faisal Shafait. Text summarization from judicial records using deep neural machines. In 2022 International Conference on Electrical, Computer, Communications and Mechatronics Engineering (ICECCME), pages 1–6. IEEE, 2022.
- [4] P. Nineesha and P. Deepalakshmi. Automated techniques on indian legal documents: A review. In 2022 Third International Conference on Intelligent Computing Instrumentation and Control Technologies (ICICT), pages 172–178. IEEE, 2022.
- [5] Sushanta Kumar. Similarity analysis of legal judgments and applying 'paragraphlink' to find similar legal judgments. Technical report, International Institute Information Technology Hyderabad India, 2014.
- [6] Ambedkar Kanapala, Sukomal Pal, and Rajendra Pamula. Text summarization from legal documents: a survey. *Artificial Intelligence Review*, 51:371–402, 2019.
- [7] Seth Polsley, Pooja Jhunjhunwala, and Ruihong Huang. Casesummarizer: A system for automated summarization of legal texts. In Proceedings of COLING 2016, the 26th International Conference on Computational Linguistics: System Demonstrations, pages 258–262, 2016.
- [8] Yisong Chen and Qing Song. News text summarization method based on bart textrank model. In 2021 IEEE 5th Advanced Information Technology, Electronic and Automation Control Conference (IAEAC), volume 5, pages 2005–2010. IEEE, 2021.
- [9] Yew-Soon Ong and Abhishek Gupta. Air 5: Five pillars of artificial intelligence research. *IEEE Transactions on Emerging Topics in Computational Intelligence*, 3(5):411–415, 2019.
- [10] Min-Yuh Day and Chao-Yu Chen. Artificial intelligence for automatic text summarization. In 2018 IEEE International Conference on Information Reuse and Integration (IRI), pages 478–484. IEEE, 2018.
- [11] P. Mahalakshmi and N. Sabiyath Fatima. Summarization of text and image captioning in information retrieval using deep learning techniques. *IEEE Access*, 10:18289–18297, 2022.
- [12] Asmaa Elsaid, Ammar Mohammed, Lamiaa Fattouh Ibrahim, and Mohammed M. Sakre. A comprehensive review of arabic text summarization. *IEEE Access*, 10:38012–38030, 2022. Title Suppressed Due to Excessive Length 13
- [13] Ravali Boorugu and G. Ramesh. A survey on nlp based text summarization for summarizing product reviews. In 2020 Second International Conference on Inventive Research in Computing Applications (ICIRCA), pages 352–356. IEEE, 2020.
- [14] B. Faizal and Sajimon Abraham. Nlp based automated business report summarization. In 2022 International Conference on Innovative Trends in Information Technology (ICITIIT), pages 1–4. IEEE, 2022.
- [15] Ayham Alomari, Norisma Idris, AznulQalid Md Sabri, and Izzat Alsmadi. Deep reinforcement and transfer learning for abstractive text summarization: A review. *Computer Speech & Language*, 71:101276, 2022.
- [16] Adhika Pramita Widyassari, Supriadi Rustad, Guruh Fajar Shidik, Edi Noersasongko, Abdul Syukur, and Affandy Affandy. Review of automatic text summarization techniques & methods. *Journal of King Saud University-Computer and Information Sciences*, 34(4):1029–1046, 2022.
- [17] Ani Nenkova and Kathleen McKeown. Mining text data. In *Mining Text Data*, pages 43–76. 2012.
- [18] Paheli Bhattacharya, Kaustubh Hiware, Subham Rajgaria, Nilay Pochhi, Kripabandhu Ghosh, and Saptarshi Ghosh. A comparative study of summarization algorithms applied to legal case judgments. In *Advances in Information Retrieval: 41st European Conference on IR Research, ECIR 2019, Cologne, Germany, April 14–18, 2019, Proceedings, Part I*, volume 41, pages 413–428. Springer International Publishing, 2019.
- [19] Amtul Waheed and Jana Shafi. Successful role of smart technology to combat covid-19. In 2020 Fourth International Conference on I-SMAC (IoT in Social, Mobile, Analytics and Cloud)(I-SMAC), pages 772–777. IEEE, 2020.
- [20] Aria Haghighi and Lucy Vanderwende. Exploring content models for multidocument summarization. In *Proceedings of Human Language Technologies: The 2009 Annual Conference of the North American Chapter of the Association for Computational Linguistics*, pages 362–370, 2009.

- [21] Aashka Trivedi, Anya Trivedi, Sourabh Varshney, Vidhey Joshipura, Rupa Mehta, and Jenish Dhanani. Extracted summary based recommendation system for indian legal documents. In 2020 11th International Conference on Computing, Communication and Networking Technologies (ICCCNT), pages 1–6. IEEE, 2020.
- [22] Ramesh Nallapati, Feifei Zhai, and Bowen Zhou. Summarunner: A recurrent neural network based sequence model for extractive summarization of documents. In Proceedings of the AAAI Conference on Artificial Intelligence, volume 31, 2017.
- [23] Shashi Narayan, Shay B. Cohen, and Mirella Lapata. Ranking sentences for extractive summarization with reinforcement learning. arXiv preprint arXiv:1802.08636, 2018.
- [24] Sebastian Gehrmann, Yuntian Deng, and Alexander M. Rush. Bottom-up abstractive summarization. arXiv preprint arXiv:1808.10792, 2018.
- [25] Jingqing Zhang, Yao Zhao, Mohammad Saleh, and Peter Liu. Pegasus: Pre-training with extracted gap-sentences for abstractive summarization. In International Conference on Machine Learning, pages 11328–11339. PMLR, 2020.
- [26] Sugam Sharma, Ritu Shandilya, and Swadesh Sharma. Predicting indian supremecourt decisions. SSRN, (3917603), 2021.
- [27] V. Gokul Pillai and Lekshmi R. Chandran. Verdict prediction for indian courts using bag of words and convolutional neural network. In 2020 Third International Conference on Smart Systems and Inventive Technology (ICSSIT), pages 676–683. IEEE, 2020.
- [28] C-C. Jay Kuo. Understanding convolutional neural networks with a mathematical model. Journal of Visual Communication and Image Representation, 41:406–413, 2016.
- [29] Reshma Sheik and S. Jaya Nirmala. Deep learning techniques for legal text summarization. In 2021 IEEE 8th Uttar Pradesh Section International Conference on Electrical, Electronics and Computer Engineering (UPCON), pages 1–5. IEEE, 2021.
- [30] Rafe Athar Shaikh, Tirath Prasad Sahu, and Veena Anand. Predicting outcomes of legal cases based on legal factors using classifiers. Procedia Computer Science, 167:2393–2402, 2020.
- [31] Wafaa S. El-Kassas, Cherif R. Salama, Ahmed A. Rafea, and Hoda K. Mohamed. Edgesumm: Graph-based framework for automatic text summarization. Information Processing & Management, 57(6):102264, 2020.
- [32] Akanksha Joshi, Eduardo Fidalgo, Enrique Alegre, and Laura Fernández-Robles. Summocoder: An unsupervised framework for extractive text summarization based on deep auto-encoders. Expert Systems with Applications, 129:200–215, 2019.
- [33] Vanyaa Gupta, Neha Bansal, and Arun Sharma. Text summarization for big data: A comprehensive survey. In International Conference on Innovative Computing and Communications: Proceedings of ICICC 2018, Volume 2, pages 503–516. Springer Singapore, 2019.
- [34] Zara Nasar, Syed Waqar Jaffry, and Muhammad Kamran Malik. Textual keyword extraction and summarization: State-of-the-art. Information Processing & Management, 56(6):102088, 2019.
- [35] Peter J. Liu, Mohammad Saleh, Etienne Pot, Ben Goodrich, Ryan Sepassi, Lukasz Kaiser, and Noam Shazeer. Generating wikipedia by summarizing long sequences. arXiv preprint arXiv:1801.10198, 2018.
- [36] Attada Venkataramana, K. Srividya, and R. Cristin. Abstractive text summarization using bart. In 2022 IEEE 2nd Mysore Sub Section International Conference (MysuruCon), pages 1–6. IEEE, 2022.
- [37] Moreno La Quatra and Luca Cagliero. Bart-it: An efficient sequence-to-sequence model for italian text summarization. Future Internet, 15(1):15, 2022.
- [38] Shusheng Xu, Xingxing Zhang, Yi Wu, and Furu Wei. Sequence level contrastive learning for text summarization. In Proceedings of the AAAI Conference on Artificial Intelligence, volume 36, pages 11556–11565, 2022.
- [39] Ishmael Obonyo, Silvia Casola, and Horacio Saggion. Exploring the limits of a base bart for multi-document summarization in the medical domain. In Proceedings of the Third Workshop on Scholarly Document Processing, pages 193–198, 2022.
- [40] Mike Lewis, Yinhan Liu, Naman Goyal, Marjan Ghazvininejad, Abdelrahman Mohamed, Omer Levy, Ves Stoyanov, and Luke Zettlemoyer. Bart: Denoising sequence-to-sequence pre-training for natural language generation, translation, and comprehension. arXiv preprint arXiv:1910.13461, 2019.
- [41] Shivam Bansal. Shivam bansal. <https://www.kaggle.com/datasets/shivamb/legal-citation-text-classification>, 2023. Last accessed: 2023/01/21.
- [42] Kanika Agrawal. Legal case summarization: An application for text summarization. In 2020 International conference on computer communication and informatics (ICCCI), pages 1–6. IEEE, 2020.
- [43] James Briggs. The ultimate performance metric in nlp. <https://towardsdatascience.com/the-ultimate-performance-metric-in-nlp-111df6c64460>, 2023. Last accessed: 2023/01/21.
- [44] Ángel Hernández-Castañeda, René Arnulfo García-Hernández, Yulia Ledeneva, and Christian Eduardo Millán-Hernández. Language-independent extractive automatic text summarization based on automatic keyword extraction. Computer Speech & Language, 71:101267, 2022.