

Domain-Specific Legal Judgment Summarizer using Latent Dirichlet Allocation

Farhana Rizvi

Preston University, Kohat,
Islamabad Campus, H-8
Islamabad

Muhammad Daud Awan, PhD

Preston University, Kohat,
Islamabad Campus, H-8
Islamabad

Malik Sikandar Hayat Khiyal, PhD

Preston University, Kohat,
Islamabad Campus, H-8
Islamabad

Amber Sarwar Hashmi, PhD

Rawalpindi Women University
Department of Computer Science

ABSTRACT

Digitalization has brought about significant opportunities and challenges for Law, IT researchers for a balanced and quality summary. A statistical and topic modeling-based strategy is presented to extract an automatic summary from the PLD for legal judgments. LDA is the measuring method to capture the most important topics, rank the summary according to the final section of the legal judgments. Summarizing legal judgments involves leveraging LDA's topic modeling to update the summarization process. To generate a quality summary using the evaluating metrics. These results show the role of the proposed algorithm in a better way the proposed algorithm is competent in computational processing and has an understandable method for implementing the PLD judgments.

General Terms

Legal Judgments, Automatic Summarization, PLD

Keywords

Digitalization, Topic modeling, Latent Dirichlet Allocation (LDA), Text summarization, Extraction-based, Summary Evaluation

1. INTRODUCTION

The internet has made it possible to access legal judgments that are complex and composed of heavy text. This automate or simplify legal Judgment and could help lawyers and other related people who manage this workload better [1]. The quickest way to summarize is to extract the necessary information from text resources. Manual summarization is expensive. It collects all related content and then summarizes it. Automatic text summarization (ATS) is an effective method. Extractive summarization is a common approach to generating summaries, where significant portions of the original text are selected and combined to form a concise summary. Automatic text summarization selects valuable sections of text and tries to generate clear summaries [2][3][4][5]. This process identifies key sentences, phrases, or words that encapsulate the ideas of the text. Extractive summarization uses techniques, such as sentence scoring, Ranked-based sentences, and top-ranked sentences. It is relatively straightforward; sentences or phrases are directly selected from the source text. Extraction-based summarization methods are favored but Abstraction-based summarization is not favored. Abstraction-based methods often rely on advanced techniques such as sequence-to-sequence models and attention mechanisms. Abstractive summaries have many challenges and

is an active area of research in NLP. TF-IDF (Term Frequency-Inverse Document Frequency) is used to identify relevant sentences for summarization like a limitation of extraction-based methods. While TF-IDF is a valuable metric for assessing the importance of words within a document, it might not capture the nuances of certain domains or texts with complex structures [6]. Abstraction-based methods can generate novel sentences that encapsulate the core ideas using different contexts. To choose the right approach based on the nature of the content, the desired summary quality, and the specific goals of the summarization task. The design of domain-specific summarization systems aims to address the limitations of general summarization methods by concentrating on specific fields or domains. With a Particular domain process for customizing the summarization, these systems can better capture the nuances of relationships in place of general methods. A domain-specific summarizer can produce more accurate and informative summaries. In this paper by integrating LDA-derived topic information into the summarization, enhancing the summarizer's ability to capture the essence of the documents in your domain. Creating domain-specific versions of metrics like ROUGE is a Recall-Oriented Understudy for Gisting Evaluation valid approach when evaluating text summarization in specialized domains such as legal texts. LDA summarization provides the potential to generate more focused and informative summaries that align with the specific themes and content of legal documents. This approach can help legal professionals quickly grasp the core aspects of complex legal cases.

2. RELATED WORK

The limited availability of large-scale training data increases the relevance of unsupervised techniques. Unsupervised summarization has been the focus of several proposed works on extractive methods, which use highly ranked candidates to create summaries instead of depending solely on word frequency [8]. Extractive and Abstractive are two Text summarization methods. Extractive method and Abstractive method [9]. Ripple-down Rules are utilized in LEXA, allowing domain experts to select sentences in real-time, consequently adding a new pattern based on their selections [10]. The catchphrase used in the analysis shows that the researchers sought to identify and extract key phrases that capture the essence of the content.[11]. Text summarization generates a summary using important information. Such as news, reviews, and scientific articles [12]. The main idea of existing summarization is to extract important sentences. An abstractive expression is to get the main idea of the

document The main idea of existing summarization is to extract salient sentences. An abstractive expression is to get the main idea of the document [13]. A summary is useful for saving time and for quickly accessing large amounts of documents. [14]. ATS serves as a solution for the increasing information overload by automating the process of condensing texts while preserving their meaningful content. However, the intricacies of language and the need to capture context, nuances, and main ideas make ATS a task that continues to push the boundaries of NLP research. Researchers are constantly developing and refining techniques to enhance the effectiveness and sophistication of ATS systems, contributing to advancements in the field as a whole [15][16]. Creating human-quality summaries is indeed a complex challenge due to the multi-faceted nature of the task while automatic summarization continues to evolve with advancements in natural language processing and machine learning techniques [17]. The evaluation of comparison between a system-generated summary and a reference summary can be conducted using both manual and automatic methods. While manual evaluation involves human assessors. Automatic evaluation methods utilize various computational techniques to measure the similarity or quality of the summaries [18].

3. PROPOSED APPROACH

The Proposed approach formulates a domain-specific Extractive summarization method using statistical measures. An extractive automatic summarization flow diagram illustrates the process and introduces topic modeling and statistical methods to enhance the information diversity of summaries. This introduces the topic of diversity in the legal manuals. (PLD, YRL, CLC).

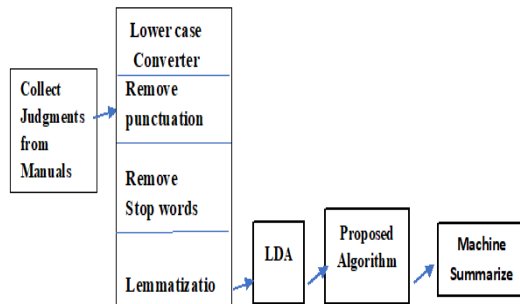


Figure1: Flow diagram

3.1 Legal Judgments

Legal Judgments are text documents collected from [19] legal Manuals (PLD). They have complex structure and often lack a clear expression of the judges' opinions. All judgments are different in their types, length, constriction, and they are organized by index and topic.

3.2 Legal Judgment Preprocessing

Legal judgments are similar to text documents, these are combinations of stop words like is, are, an, etc. these words should be removed instead of using them as topic words. A stop word list can be generated with the help of experts.

3.2.1 Convert legal judgments into lowercase

When legal judgments are converted into lowercase it shows uniformity in the judgment's words. Important step "Word" and "word" both are the same in nature.

3.2.2 Sentence Tokenization

It is the process of separating words and sentences from the given judgments. The tokenization algorithm built in the spaCY performs well (free open-source library and easy-to-manage system for extraction for NLP) since the tokenizer is trained on a data set of legal judgments [20].

3.2.3 Judgment representation and features extraction

Judgments can be represented in numerical form using different methods. Such as TF-IDF or embedding methods to transform sentences into numerical form. For example, gathering thousands of legal case reports to generate a TF*IDF matrix. Thorough which counts term frequency $TF * IDF_t = TF_t * 1 / \log N_{DF_t}$ where N refers to the number of judgments TF_t is the total counts of term, and DF_t is the number of judgments in which to appears [21].

3.2.4 Legal Judgments into Topics

An LDA generative probabilistic model is used to collect judgments by topic. It is just like probability distributions over legal terms [22]. Given a field of w distinct terms, several topics k, α , and β are smoothing parameters. And get a prior distribution over judgment lengths. LDA represents the topic using the Dirichlet distribution. Under Dirichlet distribution, these terms are independent of each other.

$$(W, Z, \Theta, \Phi | \alpha, \beta) = \prod_i P(\phi_i | \alpha) \prod_j P(z_j | \beta) \prod_{t=1}^N \prod_{k=1}^K P(w_{j,t} | \phi_{z_j, t})$$

related to the topics [23]. Each judgment is as a topic distribution and represented as:

Above said distribution, Θ presents the topic mixture, Z is the set of topic assignments, W is the words of the corpus, Φ is the topics, and α and β are hyperparameters [24].

3.3 Text Rank

Legal judgment sentences can be ranked using graph algorithms like Text rank. Text rank is an unsupervised approach and is widely used for text summarization. It generates an extraction base summary.

Flow diagram will explain the required steps.

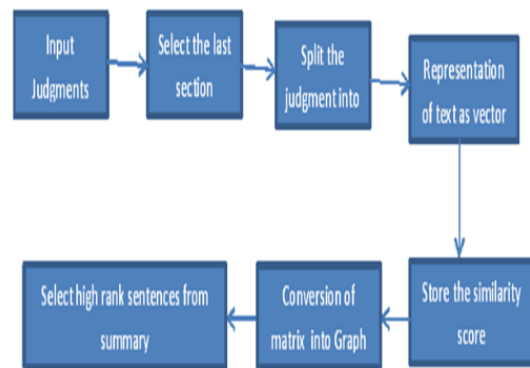


Figure2: Flow diagram of Text Rank

Input the required judgment for summary and split into sentences, representing it as a vector. Many similarity techniques determine the similarity relation between the sentences based on the overlapping content [Barrios et.al 2016]. These sentences are vertices and vertices are related to draw edges. v_i and v_j edges are used to the vertices score. Using the following formula:

$$S(V_i) = (1 - d) + d * \frac{\sum_j \text{In}(V_i)}{\text{out}(V_j)} S(V_j)$$

Here d as checker and set 0 and 1. $\text{IN}(v_i)$ set of vertices that point to it. $\text{Out}(v_j)$ set of vertices that point to it. It will find the score of each sentence (vertices) and arrange them in an order and rank and make selections of associated values of vertex [25].

3.4 Compression ratio

In the context of LDA or other topic modeling techniques, the compression ratio is adjusting various parameters, such as the number of topics, the threshold for topic inclusion, or the length of the final summary. By finding the compression ratio, It can strike a balance between reducing document length and preserving the essential legal information required for the task.

Compression ratio=

Number of sentences in a summary

Total number of sentences in a judgment

[24].

3.5 Topic Distribution (θ)

LDA generative probabilistic model It is just like probability distributions over legal terms [22]. LDA represent Topic using Dirichlet distribution under Dirichlet Distribution these terms are independent. This model provides different Topics equal to the manual describing Topic [23]. Each judgment represented as a topic distribution. It is the probability distribution of the Topic of the given document. It can be formulated as follows.

$$\theta_j^d = \text{frac} \left\{ C_{\{dj\}}^{\{DT\}} + \alpha \right\} \left\{ C_{\{dj\}}^{\{DT\}} + T\alpha \right\}$$

Here θ represent Topic distribution, $C_{\{dj\}}^{\{DT\}}$ shows the number of times a topic j is assigned to words of judgment d , T is the number of topics and α is the hypermeter [24].

3.6 Topic diversity

Latent Dirichlet Allocation (LDA) topic diversity relates to the number and variety of topics extracted from a manual. A diverse set of Topics suggests that the judgments in the manual cover a wide range of subjects, while less diversity may indicate a more focused collection. It determined the cosine similarity between the two Topics, mixture of summary and original document.

Topic Diversity = cosine similarity (topic mixture (Summary), Topic mixture (Judgment) [24].

3.7 Redundancy Rate

Redundancy rate is a quality assessment metric for text generation models and summarization algorithms. Lower redundancy is often associated with higher-quality outputs.

To calculate the redundancy rate, you typically compare the total amount of repeated or duplicated content to the Total content length. It is as a percentage or a ratio.

For example:

Redundancy Rate (%) = (Length of Redundant Content / Total Content Length) * 100 [24].

3.8 Evaluation Metric (ROUGE)

ROUGE is a Recall-Oriented Understudy for Gisting Evaluation [26] to determine the quality of the automatic summary. ROUGE-N(ROUGE01 and ROUGE-2), ROUGE-L, and ROUGE-SU4, ROUGE-N is the unit of similarity between the system summary and the reference summary based on the n-gram comparison and overlap. The given formula is applied for calculations.

ROUGE-N =

$$\frac{[\sum_s \in \{Ref. Summaries\} \sum_{gram_n \in S} Count_match(gram_n)]}{[\sum_s \in \{Ref. Summaries\} \sum_{gram_n \in S} Count(gram_n)]}$$

÷

$$[\sum_s \in \{Ref. Summaries\} \sum_{gram_n \in S} Count(gram_n)]$$

It explains that N represents the length of the N -gram count match (gram N), which is the maximum number of N -grams that occur in both the reference Summary and the candidate summary. ROUG-1 and ROUG-2 are common and determine the number of overlapping unigrams and bigrams. A Modified ROUG-SU is ROUG-SU4, and an updated version of ROUG-S[27]. ROUG-SU4 skips a maximum distance of 4 between the bigram used. ROUGE-N is a single reference. A multi References pairwise summary level can be ROUG-N. Between a candidate summary and every reference, r_i to the reference set. A final multi-reference ROUGE-N score uses the maximum pairwise summary-level ROUGE-N scores. It is the Longest Common Subsequent method and is used to evaluate the fluency of the summary. This text determines sentence-level structure similarity. S is the system summary and R is the reference summary that contains n words in ROUG-L

ROUGE-L=LCS(S, R)/ n

calculation.

LCS (S, R) represent the length of the longest subsequent of S and R [28]

3.9 Proposed Algorithm

Input: PLD Judgment (JP)

Output: Summary (Js)

Divide the Judgments into section

For each section

Begin

Convert all sentences in a section to lower case

Remove noise by removing punctuations,

Stop words and lemmatization

end

Fetch cue words from JP using LDA.

For last section in a JP

Begin

For each sentence in section

Begin Score (sentence) = score (NER's) + score (cue words) +

Similarity score (text, title) +

Position score (sentence)

end

end

JS= { }

for last section in document

begin

Sort the judgment sentences in score descending order

Sub summary= { }

For each sentence in section

begin

If sub summary == { } then

add sentence to summary

else

find similarity (current sub summary,

Current sentence)

add a sentence to sub summary if

Similarity score <= predefined value

end

Summary=JS +sub summary;

end;

Print JS

The proposed approach produces summary of source **judgments** using a combination of statistical and topic modeling. This algorithm inputs the noise-free document and highly distributed topic words from the LDA topic modeling algorithm. Separate sections into Judgments, and the last section evaluates the score of sentences by summing up the scores of named entities, the similarity between text and order, the position of the sentence, and the cue words, highly distributed words from the dominated topics of the judgment from LDA were taken as cue words to

improve the similarity of Topics between the generated summary and the source judgment. The proposed algorithm selects the sentences as follows:

1. It adds the highest-scored sentence in the sub-summary set.
2. Adding another high-scored sentence to sub summary checks the minimum similarity of that sentence in this case. This process improves the sub-summary and eliminates redundancy. In this manner, it continues adding another subsequent highly scored sentence to the sub-summary until it reaches the specified compression ratio. It generates the final summary by summing up all sub-summaries.

3.10 Data Preprocessing

PLD Judgment is the source of Data. All steps related to preprocessing at first convert letters into lower case remove punctuations and perform lemmatization. PLD judgments are in their specific form. It is better to clean data and remove the noise using different preprocessing methods.

3.11 Topic Modeling

Topic Modeling is the method used to identify hidden topics or keywords associated with PLD judgments and is used to generate candidate summaries.

3.12 Sentence Selection

Select the desired summary compression ratio to choose the candidate sentences to produce the evaluation function of candidate sentences. This step ai in the initialization of machine summary with a desired summary compression ratio.

4. RESULTS

This approach is applied using the GENSIM package to develop the LDA to fetch the cue words to rank the sentences and NLTK.SKLEARN library is used to find distance metrics. It is a combination of different types of court Judgments. Different experiments are performed using the proposed LDA-DLJS system. It can be evaluated using two different methods. At first with the help of comparison ratio, examined by altering the number of topics to 2, 4, and 6. The proposed DLJS system results in recall, precision, and F-measures for 10%. Try to sum up the given work in these sections. It compares the results using another available algorithm. To get a specific compression ratio to obtain maximum accuracy. The proposed system was evaluated using three different metrics. These are the ROUGE1 score, compression ratio, and Topic similarity. Table 1 shows various summarization methodologies using three variant cases. Three types of ROUGE cases are described, ROUGE1, ROUGE2, and ROUGEL against Precision, Recall, and Score. The proposed system is comparatively better than other systems such as bigram and long common subsequence similarity. Table 2 shows the topic similarity of the proposed algorithm at different compression ratios. It is against the Text Rank. When there is an increase in the percentage of compression ratio. A small addition in topic similarity. Both figures present graphical and numerical results, and other tables and figures are used to represent the experimental results of the proposed system. The proposed algorithm performance is best during evaluation while other algorithm such as Text Rank slightly better. It can compare the results of recall which showed high values and F1 score is different at different compression ratio. This figure shows the proposed system performance which is maintained after change in compression ratio and showed improved performance. It can be checked in Figure5 shows good topic coherence as the original document .2% compression ratio generated summary is good according to human perspective judgment while other ratios 6% and 10% are all algorithm summaries more or less similar. All this will be clear with the help of figure6. Table3 provide a list of synonyms which are related to machine to machine generated summary.

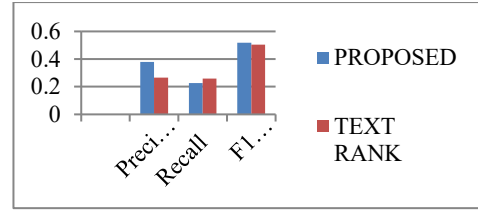


Figure3(a) results on summary quality using ROUGE1 score at a 2% compression ratio.

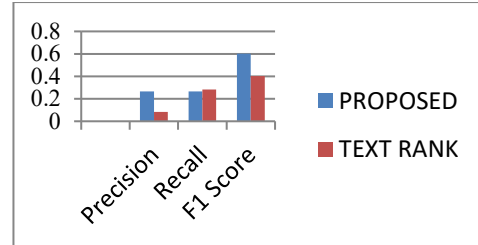


Figure3(b) results on summary quality using ROUGE2 score at a 2% compression ratio.

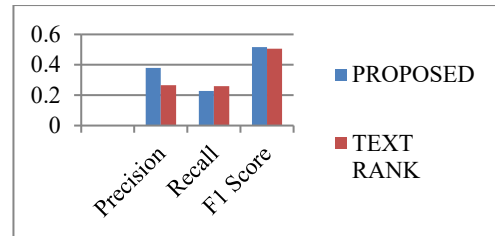


Figure3(c) results on summary quality using ROUGEL score at a 2% compression ratio.

Figure3 is showing summary quality with 2% comparison ratio in Table1.

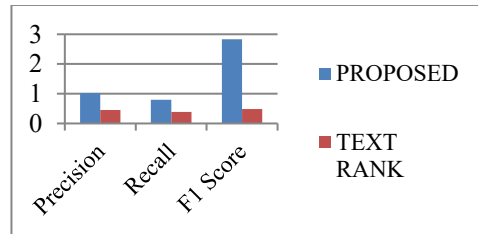


Figure4(a) results on summary quality using ROUGE1 score at a 6% compression ratio

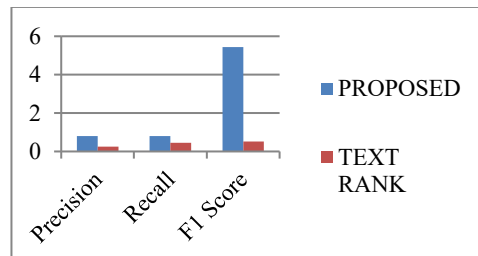


Figure4(b) results on summary quality using ROUGE2 score at a 6% compression ratio

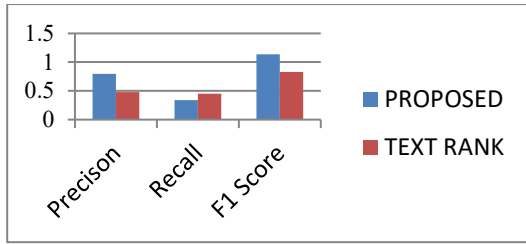


Figure4(c) results on summary quality using ROUGEL score at a 6% compression ratio

Figure4 is showing summary quality with 6% comparison ratio in Table1.

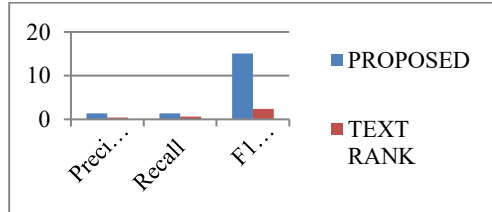


Figure5(a) results on summary quality using ROUGE1 score at a 10% compression ratio

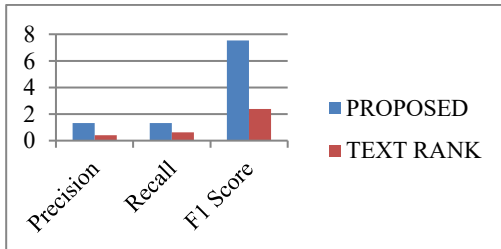


Figure5(b) results on summary quality using ROUGE2 score at a 10% compression ratio

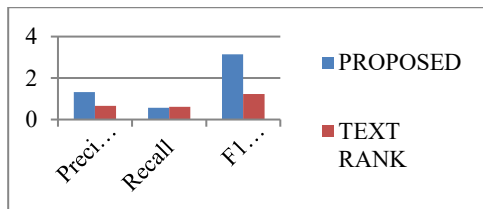


Figure5(c) results on summary quality using ROUGEL score at a 10% compression ratio

The Proposed method has a particularly higher precision than Text RANK, indicating it selects more relevant content in summaries. Both methods perform similarly, with the Proposed method slightly ahead suggesting both recover a comparable portion of relevant information. The Proposed method significantly outperforms Text RANK in F1-score a harmonic mean of precision and recall highlighting its better overall balance between identifying relevant content and avoiding irrelevant details

Figure5 is showing summary quality with 10% comparison ratio in Table1.

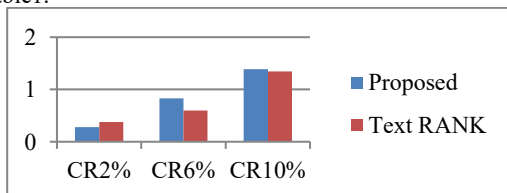


Figure6: Different compression ratios

Figure6 is showing three different ratios comparison in Table2.

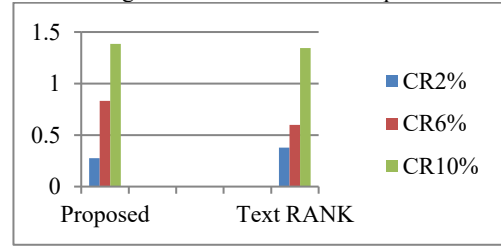


Figure7: Topic Similarity

Figure7 shows the proposed method regularly outperform text Rank at lower compression rates (2% and 6%) while at 10% both methods perform equally. The proposed method is more efficient at continuing summary quality when summarizing to shorter lengths.

Table1: Comparison Ratios

COMPRESSION RATIO 2%				
Approach	Evaluation	PRECISION	RECALL	F1 SCORE
PROPOSED	ROUGE1	0.166289	0.129336	0.153185
TEXT RANK	ROUGE1	0.282963	0.262222	0.293333
PROPOSED	ROUGE2	0.646681	0.646681	0.716908
TEXT RANK	ROUGE2	0.82963	0.282963	0.258996
PROPOSED	ROUGEL	0.646681	0.05922	0.716908
TEXT RANK	ROUGEL	0.293333	0.282963	0.532444
COMPRESSION RATIO 6%				
Approach	Evaluation	PRECISION	RECALL	F1 SCORE
PROPOSED	ROUGE1	0.498868	0.388009	1.378669
TEXT RANK	ROUGE1	0.448889	0.386667	0.48
PROPOSED	ROUGE2	1.9490043	1.9490043	6.452169
TEXT RANK	ROUGE2	0.248889	0.448889	0.730963
PROPOSED	ROUGEL	1.940043	0.17766	6.452169
TEXT RANK	ROUGEL	0.48	0.448889	0.872
COMPRESSION RATIO 10%				
Approach	Evaluation	PRECISION	RECALL	F1 SCORE
PROPOSED	ROUGE1	0.831447	0.646681	3.829635
TEXT RANK	ROUG1	0.614815	0.511111	0.666667
PROPOSED	ROUG2	3.233404	3.233404	17.92269
TEXT RANK	ROUG2	0.414815	0.614815	01.674897
PROPOSED	ROUGL	3.233404	0.296099	17.92269
TEXT RANK	ROUGL	0.666667	0.614815	1.311111

Table2: Topic Similarity

Comparison Ratio			
Approach	CR2%	CR6%	CR10%
Proposed	0.277148936	0.831447	1.385745
Text Rang	0.377777778	0.6	1.342857

5. IMPLEMENTATION

In implementation, select the PLDs for preprocessing text and extract many latent topics in the source judgments using topic modeling to generate a summary. The evaluation with precision, recall, and F1 score, and showed the best performance while existing approaches using F1 score value. The expected results of the ROUGE score tally that all preprocessing steps have been successfully completed, including noise reduction, converting text to lowercase, and removing stop words. This process achieves effective summarization without any loss of information. The performance of the proposed model is 75% better than other algorithms used for document similarity. Its highlight the superior performance of the proposed model in terms of coherence score and word associations, rather than relying on text rank. Candidate summary sentences are generated using the LDA method. This study focuses on the topic of word processing using LDA. The findings indicate that LDA is crucial for implementing legal judgments, with existing evaluations showing that its performance increases by 85% when compared to Doc2Vec. The study examines both the similarity of judgments within manual evaluations and the similarity of judgments with previously unseen cases. It is evident that topic modeling can enhance sentence selection to improve topic similarity between machine-generated texts and the original legal judgments. source text judgment. Following steps will be involved, the LDA algorithm [29][30] captures the related words to the PLD manual. It helps in replicating the source text's complex contents better for generating a summary. this study proves that topic modeling could be beneficial for sentence selection to improve the topic similarity between machine-generated and source text judgment. With this method generated summary can get on specified compression ratio. It is based on produced summaries at different compression ratios for the input text for several ROUGE metrics. The result of this study proves that the proposed system is better than existing study methods on a similar score approach. Another method Text Rank to rank sentences in legal judgment summarization is an impactable approach for identifying the most valuable and relevant sentences within a document. In this case, Text Rank extracted sentences. This is an important performance because it is involved in the internal implementation of the page rank algorithm and the generation of the similarity matrix [31]. The generated summary captures most of the sentences, scoring 0.15 for ROUGE1 and 0.13 for ROUGE2 on the F1 measure and it is considered a good summary. Given Figure 5 and table the topic similarity explains how the machine-generated summary topics are determined by calculating cosine similarity between topic distributions of summary with the source Judgment. Domain specific topic words help to rank the sentences and easily capture the topics of the source Judgments. In conclusion, when the users increase the size of the machine-generated summary, the topic similarity should increase using more topics from the source Judgments. And it summarized more topics. Present facts and figures show the worst similarity because of variation of values. Comparatively the summary generated with the proposed algorithm is more consistent for example to change in compression ratio 0.45 to 0.46 for the compression ratio % to 10%. Text rank algorithm values between the proposed algorithm. The proposed algorithm covers more topics from source judgments. Figures 3 and Table 6 show some confusion, and removing this confusion will generate the best summary.

6. CONCLUSION

This paper presents the Domain specific summarizer for legal judgments. It is based on sequence steps to generate the summary. It can be significantly improved through the

incorporation of topic modeling using Latent Dirichlet Allocation (LDA). By leveraging LDA to identify key topical structures within documents, the summarization process becomes more focused and relevant to the specific domain. This enhances the quality and contextual accuracy of the resulting summaries. Here a comparison is given between the proposed and other existing methods. The proposed LDA topic modeling algorithm that captures the main topics of the source documents. The evaluation matrices are conducted to assess the performance of the summarizers. At first, ROUGE metrics are used to get the highest F1 score with the compression ratio of 2%. In the second step, topic similarity of the machine-generated summary is calculated to cosine similarity between topic distributions of the summary with the proposed algorithm source document and maintain consistency. These values vary from 0.45 to 0.46. At the end, it is assessed manually. The grade concerned the text quality and non-redundancy in summary at a higher compression ratio. The presented results show that this model is quite successful. In the future, try to create a hybrid model that merges LDA with deep learning techniques, such as BERT. A comparison shows that LDA is capable of revealing latent topics, while BERT provides robust background comprehension and language generation abilities. The combination of both models can facilitate dynamic, context-aware summarization that remains secure in the domain's core topics.

7. REFERENCES

- [1] Moratanch and S Chitrakala. 2016. A survey on abstractive text summarization. In *Circuit, Power and Computing Technologies (ICCPCT)*, 2016 International Conference on, pages 1–7. IEEE
- [2] Chaudhary, Nidhi, and Shalini Kapoor, Key Phrase Extraction Based Multi-Document Summarization, *International Journal of Engineering Trends and Technology*.13(4) (2014) 148–53. <https://doi.org/10.14445/22315381/ijett-v13p232>.
- [3] Gambhir, Mahak and Vishal Gupta, Recent Automatic Text Summarization Techniques: A Survey, *Artificial Intelligence Review* .47 (1)(2016)166. Doi:10.1007/s10462-016-9475
- [4] Allahyari, Mehdi, Seyedamin Pouriyeh, Mehdi Assefi, Saeid Safaei, Elizabeth D, Juan B, and Krys Kochut, Text Summarization A Techniques Brief Survey, *International Journal of Advanced Computer Science and Applications*.8(10)(17). Doi:10.14569/ijacsa.2017.081052
- [5] Mekuria, Getahun Tadesse, and Aniket S Jagtap, Automatic Amharic Text Summarization Using NLP Parser, *International Journal of Engineering Trends and Technology*. 53(1) (2017) 52–58. <https://doi.org/10.14445/22315381/ijett-v53p210>
- [6] Nenkova and Kathleen McKeown. 2012. A survey of text summarization techniques. In *Mining Text Data*, pages 43–76. Springer.
- [7] Radev, Dragomir R., Simone Teufel, Horacio Saggion, Wai Lam, John Blitzer, Hong Qi Arda Çelebi, Danyu Liu and Elliott Drabek Evaluation Challenges in Large-Scale Document Summarization, *Proceedings of the 41st Annual Meeting on Association for Computational Linguistics-ACL '03*. (2003). Doi:10.3115/1075096.1075144
- [8] Nenkova, A.; and Vanderwende, L. 2005. The impact of frequency on summarization. Microsoft Research, Redmond, Washington Tech. Rep. 101.

- [9] Belwal, R. C., Rai, S., & Gupta, A. (2021). Text summarization using topic-based vector space model and semantic measure. *Information Processing & Management*, 58(3),102536. <https://doi.org/10.1016/j.ipm.2021.102536>
- [10] Filippo Galgani, Paul Compton, and Achim Hoffmann.2015Lexa: Building knowledge bases for automatic legal citation classification Expert Systems with Applications, 42(17):6391– 6407.
- [11] Filippo Galgani, Paul Compton, and Achim Hoffmann.2012b. combining different summarization techniques for legal text In Proceedings of the Workshop on Innovative Hybrid Approaches to the Processing of Textual Data, pages 115–123. Association for Computational Linguistics.
- [12] See et al. 2017; Nikolov et al. 2018; Narayan et al. 2018; Chu and Liu 2019
- [13] See et al. 2017; Narayan et al. 2018; Liu and Lapata 2019.
- [14] N. Moratanch et al “A survey on abstractive text summarization” Conference:2016International Conference on Circuit, Power and Computing Technologies (ICCPCT)
- [15] T. Cai, M. Shen, H. Peng, L. Jiang, and Q. Dai, “Improving transformer with sequential context representations for abstractive text summarization,” *Natural Language Processing and Chinese Computing*, pp.512–524, 2019
- [16] J. Wang, Y. Hou, J. Liu, Y. Cao, and C. Y. Lin, “A statistical framework for product description generation,” in *Proceedings of the Eighth International Joint Conference on Natural Language Processing*, Taiwan, China, pp. 187–192
- [17] Torres-Moreno, Juan-Manuel, Automatic Text Summarization: Some Important Concepts, *Automatic Text Summarization*. (2014) 23–52 Doi:10.1002/9781119004752. ch2.
- [18] Kiyomarsi, Farshad, Evaluation of Automatic Text Summarizations on Human Summaries, *Based, Procedia - Social and Behavioral Sciences*. 192 (2015)83–91. Doi: 10.1016/j.sbspro.2015.06.013
- [19] <http://www.keralawyer.com/asp/sub.asp?pageVal=judgement>
- [20] Natural Language Processing With spaCy in Python – Real Python Access time: 12:23:45 Am
- [21] Case Summarizer: A System for Automated Summarization of Legal Texts Seth Polsley Pooja Jhunjhunwala* Department of Computer Science and Engineering, Texas A&M University spolsley,pj861992, huangrh@cse.tamu.edu
- [22] D. M. Blei, A. Y. Ng, and M. I. Jordan. “Latent Dirichlet allocation”. *Journal of Machine Learning Research* Vol.3 (2003) 993-1022.
- [23] [ijais12-450384.pdf](#) Access time: 12:12:34 am
- [24] Extractive Summarization of Bible Data using Topic Modeling (ijettjournal.org)Access time: 10:12:34am
- [25] Gupta, Vishal and Gurpreet Singh Lehal, A Survey of Text Summarization Extractive Techniques, *Journal of Emerging Technologies in Web Intelligence*.2(3) (2010).Doi:10.4304/jetwi.2.3.258-268 .
- [26] Lin CY. Rouge: a package for automatic evaluation of summaries. In: *Text summarization branches out*;2004, Jul. p.74–81
- [27] Saziabegum S, Sajja PS. Review on text summarization evaluation methods. *Indian J Comput Sci Eng*. 2017 ;8(4):497500.
- [28] Extractive summarization of Malayalam documents using latent Dirichlet allocation: An experience (degruyter.com)
- [29] Blei, David M, Probabilistic Topic Models, *Communications of the ACM*.55(4) (2012) 77– 84 .Doi:10.1145/21]33806.2133826
- [30] Griffiths, Thomas L. and Mark Steyvers A Probabilistic Approach to Semantic Proceedings of the Twenty-Fourth Annual Conference of the Cognitive Science Society. (2019) 381-86.Doi:10.4324/9781315782379-102.
- [31] Mihalcea, Rada, Tarau and Paul TextRank:Bringing Order into Text, *Proceedings of the Conference on Empirical Methods in Natural Language Processing*. (2004) 404–11.