

How Effective are Different Machine Learning Algorithms in Predicting Legal Outcomes in South Africa?

Joe Khosa^{1,*}, Daniel Mashao², Ayorinde Olanipekun³, Charis Harley⁴

¹Department of Engineering Management, Faculty of Engineering and the Built Environment, University of Johannesburg, South Africa

²Faculty of Engineering and the Built Environment, University of Johannesburg, South Africa

^{3,4}Data Science Across Disciplines Research Group, Faculty of Engineering and the Built Environment, University of Johannesburg, South Africa

(Received: July 15, 2024; Revised: July 26, 2024; Accepted: August 14, 2024; Available online: October 19, 2024)

Abstract

This study examines the effectiveness of different machine learning algorithms in predicting legal outcomes in South Africa's Judiciary system. Considering the advancement of artificial intelligence in the legal sector, this research aims to assess the effectiveness of various machine learning algorithms within the legal domain. Text classification is done using machine learning algorithms, including Logistic Regression, Random Forest, and K-Nearest Neighbours, with datasets obtained from a state legal firm in South Africa. The datasets undergo diligent data cleansing and pre-processing methods, encompassing tokenization and lemmatization techniques. This study evaluates these models' applications through accuracy metrics. The findings demonstrate that the Logistic Regression model attained an accuracy rate of 75.05%, whereas the Random Forest algorithm achieved an accuracy rate of 75.08%. On the other hand, the K-Nearest Neighbours algorithm exhibited no optimal performance, as evidenced by its accuracy rate of 62.76%. This study provides valuable insights for legal professionals by addressing a specific research question about the successful application of machine learning in South Africa's legal sector. The results indicate the possibility of using machine learning to predict the outcomes of criminal legal cases. Additionally, this study highlights the significance of responsibly and ethically implementing machine learning within the legal field. The results of this study enhance our comprehension of the prediction of legal outcomes, establishing a foundation for future investigations in this dynamic area of study. A limitation of this study is that the data was obtained from a single law firm in South Africa.

Keywords: Legal Predictions, Legal Outcomes, Artificial Intelligence, Machine Learning Algorithms, Logistic Regression, Random Forest, K-Nearest Neighbors, Legal Document Analysis, Legal Aid South Africa

1. Introduction

In recent years, the field of Artificial Intelligence (AI) has experienced a significant improvement, with many advantages observed in its service and non-service systems. AI can replicate human functions such as pattern recognition, natural language processing, and decision-making based on historical data [1]. AI methodologies and tools can transform numerous occupations, including law, particularly on tasks that involve evaluating legal evidence, examining legally binding contracts, and analyzing legal cases. Implementing AI in the legal sector has several advantages: automation, intelligent decision-making, improved customer satisfaction, research, data analysis, resolution of complex problems, and efficient management of repetitive tasks [2], [3]. However, like any technology's implementation, this does not go without challenges. In the legal domain, challenges include overcoming language barriers and considering societal factors. Legal practitioners in South Africa are known to traditionally carry large files of papers; anything to do with technology will receive high resistance. Some challenges include reluctance to embrace new technologies, restrictions on data sharing, and substantial implementation costs associated with AI fields [4]. In South Africa, legal practitioners in governmental or private sectors often undertake cases without employing systematic methods to predict outcomes in advance. Instead, they heavily rely on precedent or personal expertise to assess the probability of success, commonly called evaluating the case's merits. However, this approach needs a more systematic

*Corresponding author: Joe Khosa (joe.khosa@gmail.com)

DOI: <https://doi.org/10.47738/jads.v5i4.215>

This is an open access article under the CC-BY license (<https://creativecommons.org/licenses/by/4.0/>).

© Authors retain all copyrights

approach. The current legal processes in South Africa have a lot of inefficiencies, such as court backlogs, high client legal fees, and significant time commitments for lawyers; by implementing a dependable prediction system, lawyers can gain data-driven insights into case probabilities, enabling them to offer more precise advice to clients and potentially mitigate these challenges [5]. This system would improve decision-making, reduce court backlogs, and optimize resource allocation within the legal system. Thus, scholars must focus on creating and applying systematic approaches to predict legal case outcomes within the South African legal domain. This research aims to overcome these challenges by utilizing machine learning algorithms and data analytics to develop a predictive model that enhances decision-making in legal practice, ultimately enhancing legal processes, boosting client satisfaction, and increasing efficiency in the legal system by providing lawyers with a systematic tool to evaluate case probabilities.

2. Literature Review

Scholars at the University of Alberta created an AI model to evaluate conflicting legal evidence and deliver legal judgments. The models can predict upcoming cases. The objective is to leverage technology to enhance the quality of human legal decision-making. Scholars have researched prediction models within the legal domain, emphasizing text classification models, which are common in the legal domain. The purpose of these models is to assist lawyers in saving time by automating the process of evaluating voluminous legal documents. Heike researched the prediction of court case outcomes [4], whereas Chandra focused their investigation on charge prediction using neural networks [6]. The potential for AI to transform the legal industry is significant. However, its adoption on a large scale is dependent upon the development of technologies that are transparent, interpretable, and understandable to both legal professionals and the general public [7].

Khosa emphasizes how inconsistencies in legal case outcomes and court backlogs negatively affect the criminal justice system's effectiveness [8]. Implementing outcome-predictive models is crucial for improving the South African criminal justice system. Legal practitioners and policymakers can gain valuable insights into the factors influencing judicial decisions and the likely outcomes of court cases using machine learning algorithms and statistical models. This proactive approach represents a significant step towards strengthening the integrity and efficiency of the legal framework.

Oladoyinbo et al. argue that AI technology's advancement is impeded by ethical and reliability concerns, underscoring the necessity for continuous examination and enhancement of the legislation. This is associated with legal practice, which demands high problem-solving abilities and emotional intelligence, skills that machines cannot replicate [9]. AI has played a crucial role in facilitating the administration of justice during the COVID-19 pandemic, leading to significant changes in the justice system [10]. According to Chandra et al. [6], AI has the potential to aid human judges in expediting decision-making processes. Machine learning has primarily supported legal decision-making in criminal detection, finance, and sentencing. The use of AI in law has raised concerns that the inappropriate use of AI may lead to biased decision-making [5].

Scholars have compared the performance of machine learning against human experts in various industries, including the legal domain. Menezes et al. have investigated the outcome of trials in the Brazilian judiciary system; this was informed by the challenge faced by the Brazilian judiciary system of a high number of cases received each year. To predict their outcomes, they trained three deep learning architectures, ULMFiT, BERT, and Big Bird, on 612,961 Federal Small Claims Courts appeals within the Brazilian 5th Regional Federal Court. They compared the predictive performance of the models to the predictions of 22 highly skilled experts. They found that all models outperformed human experts, with the best one achieving a Matthews Correlation Coefficient of 0.3688 compared to 0.1253 from the human experts. Their study demonstrates that natural language processing and machine learning techniques provide a promising approach to predicting legal outcomes. It can be deduced from their study that it is possible to use deep learning models to predict outcomes of appeals in Brazilian courts, achieving performance that is better than that resulting from analysis by human experts [11].

In another study, Orosz et al. experimented with the performance of human legal editors and a machine learning model to ascertain whether machine learning could reach human performance and whether machine learning models could replace human editors. During the experiment, the performance of three competence groups was examined: legal

editors, lawyers and laymen. Every group was divided into two parts. The first group could use the results of the machine learning algorithm as assistance. The second group completed the labelling without assistance. The results showed that the proposed machine learning solution, which found 48% of the information in the reference dataset, significantly outperformed the average of the legal editors in the whole test. An observation was made wherein those participants who used the computer assistance were slower, but their precision increased by more than 50% [12]. Mumcuoğlu et al. investigated the prediction of court rulings in Turkey's higher courts. Their study found that higher court rulings in Turkey can be predicted with reasonable accuracy, as shown by considering several alternative measures. Direct comparison with previous work could not be made because earlier systems were developed for other languages and very different legal systems [13].

In summary, AI's impact on the legal field is promising, but its implementation requires the creation of user-friendly and transparent technologies. Implementing legal decision-support systems can enhance access to justice, administrative efficiency, and judicial consistency. However, policymakers and legal professionals must exercise caution and ensure these technologies are developed responsibly and ethically. In this study, we have utilized a dataset obtained from a legal firm in South Africa that encompasses data about resolved criminal court cases, intending to forecast criminal case verdicts.

The study employed a systematic four-stage approach, encompassing data collection, data cleansing, application of machine learning algorithms, and model assessment. Apart from the introductory section, the rest of this document is structured into four primary segments. In Section 2, the methods and materials employed for the study are presented. Section 3 delineates a range of algorithms, while Section 4 comprises a discourse on the investigation and a synopsis of prospective future research.

3. Research Methodology - Materials and Methods

The methodological approach entails a sequential process, beginning with data collection, followed by data pre-processing, which encompasses data exploration and cleaning and culminates in executing machine learning algorithms. This four-step methodology is illustrated in figure 1 below.

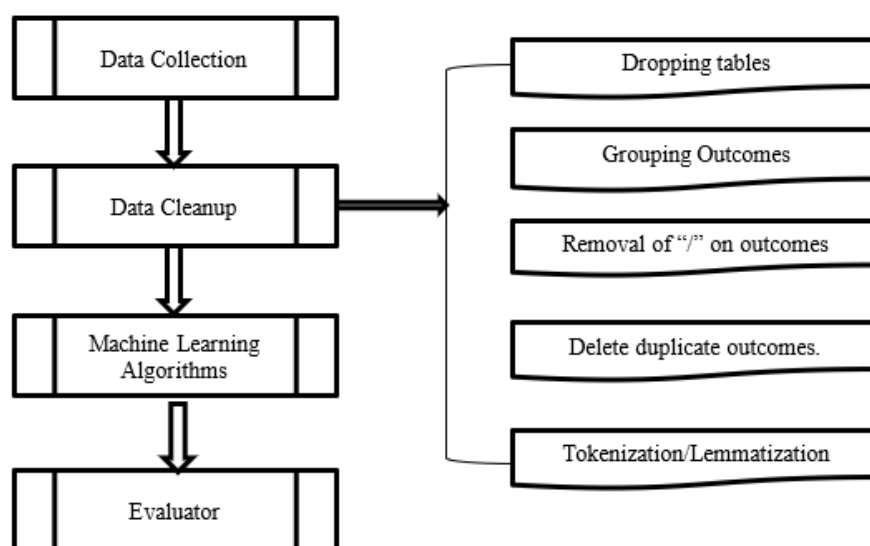


Figure 1. Workflow of Criminal court cases analysis using ML and NLP

3.1. Data Collection

The dataset was acquired from a legal firm operating under the auspices of a government parastatal in South Africa. It was collected from the electronic records of criminal court cases between 2015 and 2021. The firm employs a system known as electronic Legal Aid Administration (eLAA) to record and manage its cases [14]. The dataset contains consolidative information about the accused, charges, court proceedings, pleas, and verdicts. Furthermore, it provides detailed information about lawyers engaged in each case, such as their date of admission and professional designation.

The dataset may be used in another research. Identifying information such as names and identification numbers was redacted to safeguard the clients' anonymity. Demographic information about the accused, including gender and race, was determined to be irrelevant to the problem at hand and thus was not incorporated into the final model. [Table 1](#) below illustrates the features used in the model.

Table 1. Model Features

Features	Description
InstructionNumber	Unique identifier of a matter appearing before a criminal court.
Court Type	Different levels of courts, Lower court, Higher court or Supreme Court of Appeal
Practitioner_type	Internal Lawyers or Outsourced Lawyers.
Applicant_Outcome	Outcome(verdict) of the matter in a criminal trial.
Criminal Charges	Criminal charges against the accused.

A legal case can result in various court rulings, such as fines, suspended sentences, withdrawals, and other possible outcomes. The chosen outcomes for this study are limited to instances of imprisonment and acquittal. The focus is on the adverse outcomes (Mandate terminated, Imprisonment 3 months, Imprisonment 3 months – 1 year, Imprisonment 1 – 5 years, and Imprisonment > 5 years) and the positive outcome being acquittal (Struck off the roll, Not Guilty and Withdrawn by State). Adverse outcomes imply that the accused will not have triumphed in the legal proceedings, leading to their incarceration. Conversely, positive outcomes indicate that the defendant will be acquitted and released from custody. The outcome grouping is detailed below in [table 2](#).

Table 2. Court Outcome Grouping

Number	Outcome	Grouping
1	Struck off the roll,	Positive Outcome
2	Not Guilty,	Positive Outcome
3	Withdrawn by State	Positive Outcome
4	Imprisonment < 3month	Negative Outcome
5	Imprisonment 3 months – 1 year,	Negative Outcome
6	Imprisonment 1 year – 5 years	Negative Outcome
7	Imprisonment > 5	Negative Outcome
8	Mandate Terminated	Negative Outcome

Outcomes 1, 2, and 3 have been categorized as positive outcomes because the dismissal of a case signifies a victory for the clients being represented, and the decision to withdraw by the State is similarly seen as a favorable result. On the contrary, outcomes 4 to 8 are deemed unfavorable since they lead to the clients' imprisonment in prisons. In South Africa, the spectrum of criminal charges includes a variety of offences, spanning from theft, which is classified as a relatively minor violation, to murder, which is regarded as a grave crime. The South African Judiciary encompasses a wide range of criminal charges, exceeding 100. To facilitate our analysis, we have chosen to focus on the charges of Armed Robbery, Corruption, Murder/Attempted Murder, Rape, and Robbery, since these offences are widely seen in South Africa. These features are used as inputs in the model, as described in [table 2](#) above.

3.2. Data pre-processing

The dataset presented several factors that required cleaning, including duplicated court outcomes for the primary and co-accused accused. Specific outcomes were found to have synonyms, such as "Rape/Attempted Rape. "The term "rape" was removed from these results to facilitate more effective data processing. It was necessary to eliminate the forward slash ("/") character to facilitate data comprehension. To achieve this objective, we employed the techniques of tokenization and lemmatization, which are defined as follows.

Figure 2 below is a graphical depiction of the process followed in applying the data processing. Natural Language Processing (NLP) receives the user's query as input data. In the Pre-processing phase, we have both tokenization and lemmatization. To begin this process, we decompose the text into smaller single sentences and semantic units by tokenization. Words are then standardized using lemmatization. We then remove stop words, i.e., prepositions and articles such as “/”, “ ”, “-ing”, “-ly”, and “...” as they do not contain unique information. Once pre-processing has been completed, NLP converts the text into a language a bot can understand. The NLP process uses the BoW technique to extract features from the input sentence by computing the frequency of each word in the input text, which is later used by the neural network. Figure 2 illustrates the NLP process in the proposed system [15].

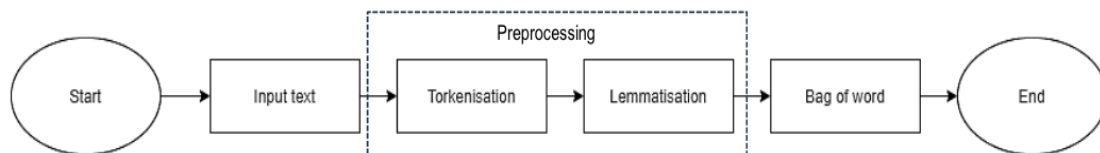


Figure 2. NLP Process

3.2.1. Tokenization

According to Devrapalli et al. [16], tokenization within NLP is commonly understood as dividing a continuous sequence of characters into individual units of meaning, typically words. Frequently, it is correlated with processes at either the lower or higher level. Despite the common tendency to categorize both tasks as "pre-processing," it is essential to note that tokenization is distinct from initial "cleaning procedures," such as eliminating extraneous tags, removing non-textual elements, and excluding elements that do not pertain to natural languages, such as mathematical or chemical formulas and programs.

3.2.2. Lemmatization

Lemmatization refers to consolidating a word's various inflected forms into a unified entity, commonly known as the word's lemma or vocabulary form. This process involves linking text containing similar meanings to a single word. Identifying the lemma of a word, which pertains to the root word as opposed to the root stem, is commonly referred to as an algorithmic technique. The interpretation of a word is contingent upon the semantic content it is attempting to express [17].

3.2.3. Featurization

The CountVectorizer methodology converts a corpus of written texts into a matrix that represents the frequency of occurrence of each word or token. Furthermore, it facilitates the pre-processing of textual data before constructing it into matrix form. This textual feature presentation architecture is highly adaptable due to its versatility. This approach was chosen because of the critical term's frequent recurrence. The frequency of a word's occurrence indicates that the word is relevant to the subject matter. The word with a higher frequency indicates that the word is more important. The count vectorizer from Scikit-learn was utilized in our study [18].

3.3. Machine Learning Algorithms

Three machine learning algorithms, the logistic regression, the random forest, and the K-nearest neighbors have been selected to create the prediction model. Random Forest has been chosen as the ideal algorithm because it can generate decision trees randomly and independently from the sampled dataset. It uses large numbers, so it does not overfit, which is a good feature for prediction [19]. The K-Nearest Neighbour algorithm uses the feature similarity to predict the value of the latest data point [20]. Furthermore, Logistic regression is considered highly interpretable because it is a linear model, resulting in its output being directly interpreted.

In comparison with Logistic Regression, Random Forest is inherently complex. Although Random Forest is considered more accurate than Logistic Regression, its complexity can be a drawback in a legal context. KNN is considered less interpretable than Random Forest and Logistic Regression [21]. We believe that the combination of these algorithms will provide a balanced view of the model. The Algorithms are further defined in detail below.

3.3.1. Logistic Regression

Logistic regression is used to test the relationship between dependent and independent variables. It can bring about distinct predictions to produce binary outcomes and predictions for continuous variables [10]. The independent and dependent variables are typically used as input and output variables. The probability is calculated using the logistic function. A logistic function is a mathematical function that uses an S-shaped curve as an input and generates an output that usually ranges between 0 and 1[22].

$$l = \log_b \frac{p}{1-p} = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots \beta_n x_n. \quad (1)$$

The intercept of the regression line is denoted by β_0 , while the estimated probability is represented by p . The coefficient of the predictor variable x_i is denoted by β_i , and the predictor variable for the coefficient β_i is x_i , where $i = 1, \dots, n$. The hyperparameter tuning process involved the utilization of the grid search method. The models and Grid search parameters used for the logistic regression in Python programming implementation are defined below. The Logistic Regression and various parameter configurations are investigated. The solvers that were tested are 'liblinear', 'lbfgs', and 'newton-cg'. The L2 regularization is the penalty that is applied. Furthermore, the model is subjected to multiple values for the inverse regularization strength (C) that range from 100 to 0.001.

3.3.2. Random Forest

The Random Forest algorithm is a well-known supervised machine learning technique widely utilized for addressing classification and regression tasks within machine learning. It is widely understood that a forest is composed of many trees and that its resilience is positively correlated with the number of trees present [23]. A random forest algorithm is a classification technique that uses data presented to it to produce multiple decision trees. Its accuracy and problem-solving ability increase proportionally with the number of trees it contains. To increase the model accuracy, averaging techniques are used on the predicted results [24]. The figure 3 below is a graphical depiction of the random forest algorithm from the sample dataset.

After applying the grid search for the hyperparameter tuning optimization using max_depth: [8,10,12,14], max_features: [60,70,80,90,100], min_samples_leaf: [2, 3, 4], and n_estimators: [100, 200, 300], the best hyperparameters for Random forest Classifier are 0.881, and grid_search: {'CV': 3, 'n_jobs': '-1', 'verbose': '2'}.

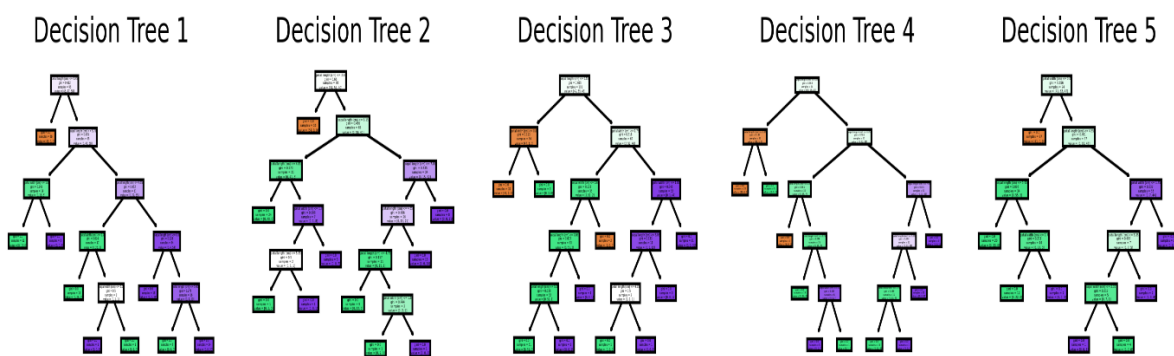


Figure 3. Random Forests Decision Tree

3.3.3. KNeighbors Classifier

According to [25], the family of the kNN algorithm includes instance-based, competitive learning, and lazy learning algorithms. The KNN algorithm is an extreme form of instance-based method because all training observations are retained as part of the model. According to [26], the kNN algorithm uses distance metric techniques to determine the similarity and dissimilarity between the data items. With the above information on the kNN algorithm, the architecture of kNN will be helpful in predictive analysis. When the unseen data instance needs to be predicted, the kNN algorithm will find the k-Most similar instances in the training dataset.

3.4. Machine Learning Ethical Concerns

Implementing machine learning raises several ethical concerns. In this section, we discuss those associated with the legal domain in the context of our model and provide strategies for mitigating them to promote fair and moral algorithm design. One of the challenges is bias; we aim to implement fairness-aware machine learning. The aim is to lessen bias by including fairness constraints in the training and evaluation procedures. Reweighting, adversarial training, and resampling are some strategies used to overcome these challenges. This paper lays the groundwork for researchers, practitioners, and policymakers to forward the cause of ethical and fair machine learning through concerted effort [27].

4. Results and Discussion

This section elucidates the results obtained from the applied algorithms on the provided dataset. We commence by elaborating on both the Unigram and Bigrams, followed by insights into the Latent Dirichlet Allocation (LDA) model. Subsequently, we delve into the actual performance metrics of these models.

4.1. Unigram and Bigrams

Figure 4(a) and figure 4(b) below illustrate the most often used word in the dataset for the criminal charge dataset. This task was achieved using Unigram and Bigrams, which explicitly emphasized the criminal charge data frame. We analyze the frequently encountered words as single units (unigrams) and in pairs (bigrams). The first stage included the preparation of the data for analysis. A data frame containing matter-outcome information was used. The text corpus was segmented into individual words or tokens to do a frequency analysis of words. The corpus was tokenized using the Natural Language Toolkit (NLTK) library. The list variable contained the generated list of tokens. In the context of frequency analysis and visualization, our first objective was to calculate the occurrence frequency of unigrams and bigrams inside the tokenized corpus. Subsequently, we produced bar plots to visualize the data that we obtained. We used the FreqDist function from the Python library NLTK to calculate the frequency distribution of the tokens.

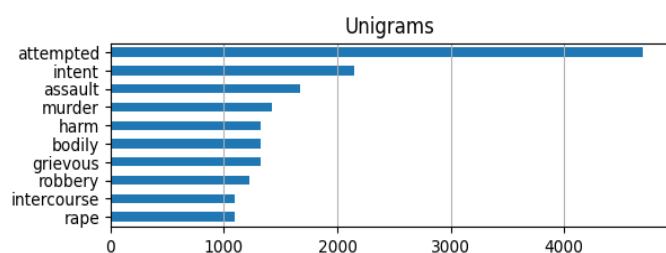


Figure 4a. Unigram: Most Frequently Used Words.

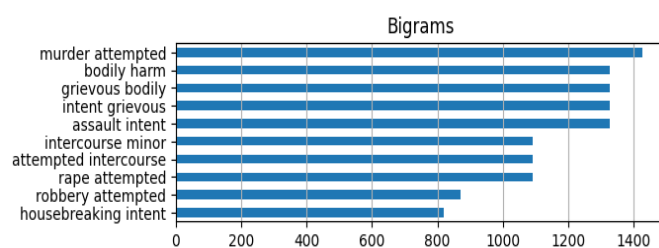


Figure 4b. Bigram: Most Frequently Used Words.

The frequency distribution of a unigram was calculated by applying the FreqDist function to the list below. The above frequency distribution was converted into a data frame having two columns named "Word" and "Freq." "Word" denotes the unigram, whereas "Freq" denotes the frequency of the unigram. We took the frequent ten individual words from the text and produced a horizontal zig-zag bar plot. The Freq is on the x-axis, while the unigram labels are on the y-axis. The above graph is illustrated in figure 4. FreqDist was also used to determine the frequency distribution of bigrams, similar to how it was done for unigrams. Figure 4(b) illustrates a horizontal bar plot displaying the top 10 common bigrams. The frequency of the bigrams is indicated on the x-axis, whereas labels of the bigrams are displayed on the y-axis.

Based on the visual representation presented in figure 4a and figure 4b, an analysis of the top 10 frequently employed words within our dataset reveals that instances of attempted murder exhibit a higher prevalence within the court charges under scrutiny. This observation is made in light of the substantial number of criminal charges, exceeding 100, documented in South Africa. This finding is consistent with the crime figures published by the South African Police Services on their official website for the period spanning from 2019 to 2023.

4.2. Latent Dirichlet Allocation

The use of LDA for modelling criminal charges, as seen in figure 5, entails using this statistical topic modelling method on a collection of textual data. The corpus is produced with great attention to detail, and this diligence is carried over to the training of the LDA model. Following that, the main charges brought against the accused are carefully visualized.

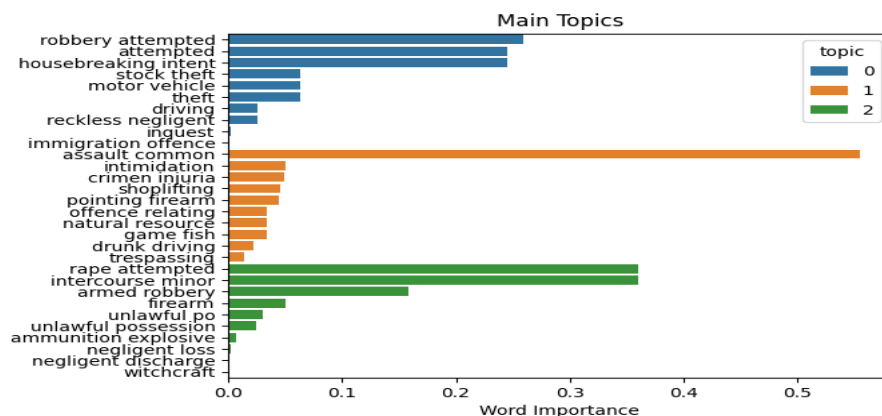


Figure 5. Criminal charge Modelling using LDA

The presence of corruption and attempted murder within the dataset mentioned above is apparent. The dataset under examination demonstrates a notable frequency of crimes classified as common assault, rape, and attempted robbery, suggesting a significant occurrence of these criminal activities. The assertion is also consistent with the crime statistics in South Africa.

4.3. Performance of the Predictive Model

Table 3 below illustrates the model performance results. The Random Forest model demonstrates a training set accuracy of 75.08% and a test set accuracy of 75.06%. The proximity of the training and test accuracies suggests that the model is not exhibiting overfitting. The K-Nearest Neighbors (KNN) algorithm attains a comparable level of accuracy to that of the Random Forest model, exhibiting a training set accuracy of 62.76% and a test set accuracy of 62.82%. The comparable training and test accuracies indicate that the model is not exhibiting overfitting. The Logistic Regression model exhibits an accuracy of 75.05% on the training set. The accuracy on the test set demonstrates an improvement to 75.12%, which is consistent with the accuracy of the other Random Forest and K-Nearest Neighbors models.

Table 3. ML Algorithm performance

Measure	Random Forest	KNeighbors	Logistic Regression
Training	75,08%	62,76%	75,05%
Test	75,06%	62,82%	75,12%

The observed performance is otherwise suboptimal, which may mean that the algorithm's efficacy may be enhanced or, conversely, there may be poor patterns in the data that the Random Forest model should identify faster [9]. The accuracy of the K-Nearest Neighbors implies that it is not the ideal method for the dataset considered or that the features do not have a substantial degree of association with the target variable. All algorithms' performance might improve with training; however, results show that Logistic Regression and Random Forest algorithms produce better outcomes than the KNeighbors.

Medvedeva et al. investigate how natural language processing tools can be used to analyze texts of court proceedings in order to automatically predict judicial decisions. They achieved an average accuracy of 75% in predicting the violation of 9 European Convention on Human Rights articles. They demonstrated that they could achieve a classification performance (average accuracy of 65%) when predicting outcomes based only on the surnames of the judges who tried the case [28]. This is below our best-performing model, which achieved an accuracy of 75.08%. In another study, Rosili et al. determined and analyzed the machine learning methods used in predicting court decisions; they achieved an accuracy rate of 70%, which is 5.08% below our best performance results [9]. Another prediction model based on contiguous word sequences (i.e., N-grams and topics) has up to a 79% accuracy rate and is now applied to cases in the European Court of Human Rights [29]. Although we could deduce from this performance that their model performed slightly better than ours, we believe that our model can still do better with more training.

5. Conclusion

We examined a dataset containing legal matters from a parastatal legal firm in South Africa, aiming to assess factors influencing a legal outcome. Machine learning algorithms such as Logistic Regression, Random Forest, and K-Nearest Neighbours were used to test the data. The findings emphasize the effectiveness of machine learning in accurately predicting legal outcomes, thereby potentially imposing a significant effect on the legal sector. Legal professionals can employ this technology to make informed decisions in specific cases, ensuring fair and just outcomes by utilizing readily available data. Several advantages are associated with applying machine learning in the Legal domain. These advantages range from quicker decision-making processes to customer satisfaction and management of repetitive tasks. The study examines the possibility of predicting the legal outcomes; this has the potential to transform the legal domain. The predicted model will be able to improve the court processes by far; court backlogs may be reduced significantly, with clients having the ability to make decisions in the conform of their homes as to which cases to pursue and which ones to let go. Most long-lasting cases result from both the defendant and plentiful being persistent on cases that may not have merits. A system that can determine beforehand whether a case has the prospect of suspects may even save clients legal fees.

This technological advancement in the legal domain will benefit clients and legal practitioners. Knowing which cases to focus their resources on will unboundedly improve efficiencies in the South African judiciary system. Several scholars have discussed the limitation of AI in replacing problem-solving abilities and emotional intelligence, which are skills that human judges and lawyers possess. South African Law Society, the South African Bar Council and the Judiciary may utilize these findings to formulate strategies and policies to improve the quality of the South African Judicial systems, particularly considering the factors that impact the legal outcomes outlined in this paper. There are still areas for further research and investigation, such as ethical concerns that come with the implementation of AI. Artificial Intelligence's responsible and ethical implementation is important to enhance access to justice and promote consistency in the South African Judicial systems. This study contributes significantly by providing a relationship between machine learning algorithms and the legal domain. A limitation of this study is that the data was obtained from a single law firm in South Africa. However, the law firm is a parastatal and employs the highest number of lawyers in South Africa. This may present a challenge where we cannot generalize our findings. Another limitation is the relatively high bias of the model against the indigenous population, which is a group that this Law firm represents. Future studies will focus on debiasing the model by examining an increased population.

6. Declarations

6.1. Author Contributions

Conceptualization: J.K., D.M., A.O., and C.H.; Methodology: J.K.; Software: A.O.; Validation: D.M. and D.M.; Formal Analysis: J.K. and A.O.; Investigation: J.K.; Resources: C.H.; Data Curation: A.O.; Writing Original Draft Preparation: J.K. and D.M.; Writing Review and Editing: A.O. and C.H.; Visualization: J.K. and D.M.; All authors have read and agreed to the published version of the manuscript.

6.2. Data Availability Statement

The data presented in this study are available on request from the corresponding author.

6.3. Funding

The authors received no financial support for the research, authorship, and/or publication of this article.

6.4. Institutional Review Board Statement

Not applicable.

6.5. Informed Consent Statement

Not applicable.

6.6. Declaration of Competing Interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

References

- [1] M. Verma, "Artificial Intelligence Role in Modern Science: Aims, Merits, Risks and Its Applications," *International Journal of Trend in Scientific Research and Development (IJTSRD)*, vol. 7, no. 5, pp. 335–342, Sep. 2023.
- [2] N. Rane, Saurabh Choudhary, and Jayesh Rane, "Artificial Intelligence-Driven Corporate Finance: Enhancing Efficiency and Decision-Making Through Machine Learning, Natural Language Processing, and Robotic Process Automation in Corporate Governance and Sustainability," *Social Science Research Network*, vol. 5, no. 2, pp. 1–22, 2024, doi: 10.2139/ssrn.4720591.
- [3] Y. K. Dwivedi, L. Hughes, E. Ismagilova, G. Aarts, C. Coombs, T. Crick, Y. Duan, R. Dwivedi, J. Edwards, A. Eirug, V. Galanos, P. V. Ilavarasan, M. Janssen, P. Jones, A. K. Kar, H. Kizgin, B. Kronemann, B. Lal, B. Lucini, R. Medaglia, K. Le Meunier-FitzHugh, L. C. Le Meunier-FitzHugh, S. Misra, E. Mogaji, S. K. Sharma, J. B. Singh, V. Raghavan, R. Raman, N. P. Rana, S. Samothrakakis, J. Spencer, K. Tamilmani, A. Tubadji, P. Walton, and M. D. Williams., "Artificial Intelligence (AI): Multidisciplinary perspectives on emerging challenges, opportunities, and agenda for research, practice and policy," *International Journal of Information Management*, vol. 57, no. 1, pp. 1–49, Apr. 2021, doi: 10.1016/j.ijinfomgt.2019.08.002.
- [4] Rinku and Gurjeet Singh, "Artificial intelligence in sustainable energy industry: status quo, challenges, and opportunities," *EPRA international journal of multidisciplinary research*, vol. 9, no. 5, pp. 234–237, May 2023, doi: 10.36713/epra13323.
- [5] J. Zeleznikow, "The benefits and dangers of using machine learning to support making legal predictions," *WIREs Data Mining and Knowledge Discovery*, vol. 13, no. 4, pp. 1–21, 2023, doi: 10.1002/widm.1505.
- [6] Geetanjali Chandra, G. R. Chandra, R. Gupta, Ruchika Gupta, N. Agarwal, and Nidhi Agarwal, "Role of Artificial Intelligence in Transforming the Justice Delivery System in COVID-19 Pandemic," *International Journal on Emerging Technologies*, vol. 11, no. 3, pp. 344–350, 2020.
- [7] H. Felzmann, E. Fosch-Villaronga, and C. Lutz, "Towards Transparency by Design for Artificial Intelligence," *Science and Engineering Ethics*, vol. 26, no. 6, pp. 3333–3361, 2020, doi: 10.1007/s11948-020-00276-4.
- [8] J. Khosa, D. Mashao, D. Ighravwe, and C. Harley, "An Examination of the Impact of a Criminal Case's Jurisdiction on the Prison Term in South Africa," *International Journal of Management Science and Business Administration*, vol. 9, no. 5, pp. 24–35, 2023, doi: 10.18775/ijmsba.1849-5664-5419.2014.95.1002.
- [9] Rosili, Rohayanti Hassan, N. Zakaria, S. Kasim, F. Rose, and T. Sutikno, "A systematic literature review of machine learning methods in predicting court decisions," *IAES International Journal of Artificial Intelligence (IJ-AI)*, vol. 8, no. 1, pp. 31–42, Mar. 2021, doi: 10.11591/ijai.v10.i4.pp1091-1102.
- [10] P. Bhilare, N. Parab, N. Soni, and B. Thakur, "Predicting Outcome of Judicial Cases and Analysis using Machine Learning," *International Research Journal of Engineering and Technology (IRJET)*, vol. 06, no. 02, pp. 1–5, 2019.
- [11] E. J. de Menezes Neto and Marco Bruno Miranda Clementino, "Using deep learning to predict outcomes of legal appeals better than human experts: A study with data from Brazilian federal courts," *PLOS ONE*, vol. 17, no. 7, pp. e0272287–e0272287, Jul. 2022.
- [12] T. Orosz, R. Vági, G. M. Csányi, D. Nagy, I. Üveges, J. P. Vadász, and A. Megyeri, "Evaluating Human versus Machine Learning Performance in a LegalTech Problem," *Applied Sciences*, vol. 12, no. 1, pp. 297–297, Dec. 2021, doi: 10.3390/app12010297.
- [13] E. Mumcuoğlu, C. E. Öztürk, H. M. Ozaktas, and A. Koç, "Natural language processing in law: Prediction of outcomes in the higher courts of Turkey," *Information Processing and Management*, vol. 58, no. 5, pp. 1–16, Sep. 2021, doi: 10.1016/j.ipm.2021.102684.
- [14] J. Makokoane, D. Khosa, and B. Obrenovic, "Applying UTAUT and Fuzzy Dematel Methods: A New Legal Aid Administration System," *the international journal of management science and business administration*, vol. 8, no. 1, pp. 24–36, Nov. 2021, doi: 10.18775/ijmsba.1849-5664-5419.2014.81.1002.
- [15] Tuan-Jun Goh, Lee-Ying Chong, Siew-Chin Chong, and Pey-Yun Goh, "A Campus-based Chatbot System using Natural Language Processing and Neural Network," *Journal of Informatics and Web Engineering*, vol. 3, no. 1, pp. 1–21, Feb. 2024, doi: 10.33093/jiwe.2024.3.1.7.

-
- [16] Dr. D. Devrapalli, S. P. Guduri, P. J. Madhuri, S. N. V. Reddy, and P. Pasupuleti, "Effective Text Processing utilizing NLP," *Indian Journal of Artificial Intelligence and Neural Networking (IJAINN)*, vol. 2, no. 1, pp. 1–7, Dec. 2021, doi: 10.54105/ijainn.B3873.122121.
- [17] D. Khyani, B. S. Siddhartha, N. M. Niveditha, and B. M. Divya, "An interpretation of lemmatization and stemming in natural language processing," *Journal of University of Shanghai for Science and Technology*, vol. 22, no. 10, pp. 350–357, 2021.
- [18] S. Khan, M. Anwar, H. Qayyum, F. Ali, and M. Nawaz, "Fake News Classification using Machine Learning: Count Vectorizer and Support Vector Machine," *Journal of Computing and Biomedical Informatics*, vol. 4, no. 01, Art. no. 01, Dec. 2022, doi: 10.56979/401/2022/85.
- [19] M. A. Afrianto and M. Wasesa, "Booking Prediction Models for Peer-to-peer Accommodation Listings using Logistics Regression, Decision Tree, K- Nearest Neighbor, and Random Forest Classifiers," *Journal of Information Systems Engineering and Business Intelligence*, vol. 6, no. 2, pp. 123–132, Oct. 2020, doi: 10.20473/jisebi.6.2.123-132.
- [20] P. Dewani, M. Sippy, G. Punjabi, and A. Hatekar, "Credit Scoring : A Comparison between Random Forest Classifier and K- Nearest Neighbours for Credit Defaulters Prediction," *International Research Journal of Engineering and Technology (IRJET)*, vol. 07, no. 10, pp. 1–6, 2020.
- [21] M. Saberioon, P. Císař, L. Labbé, P. Souček, P. Pelissier, and T. Kerneis, "Comparative Performance Analysis of Support Vector Machine, Random Forest, Logistic Regression and k-Nearest Neighbours in Rainbow Trout (*Oncorhynchus Mykiss*) Classification Using Image-Based Features.," *Sensors*, vol. 18, no. 4, pp. 1–15, Mar. 2018, doi: 10.3390/s18041027.
- [22] Shweta Agrawal, Sanjiv Kumar Jain, Shruti Sharma, and Ajay Khatri, "COVID-19 Public Opinion: A Twitter Healthcare Data Processing Using Machine Learning Methodologies," *International Journal of Environmental Research and Public Health*, vol. 20, no. 432, pp. 1–17, 2022, doi: 10.3390/ijerph20010432.
- [23] B. Charbuty, Adnan Abdulazeez, Adnan Abdulazeez, and A. M. Abdulazeez, "Classification Based on Decision Tree Algorithm for Machine Learning," *Journal of Applied Science and Technology Trends*, vol. 2, no. 1, pp. 20–28, 2021, doi: 10.38094/jastt20165.
- [24] Nasiba Mahdi Abdulkareem, N. M. Abdulkareem, Adnan Mohsin Abdulazeez, Adnan Mohsin Abdulazeez, Adnan Mohsin Abdulazeez, and A. M. Abdulazeez, "Machine Learning Classification Based on Radom Forest Algorithm: A Review," vol. 5, no. 2, pp. 128–142, 2021, doi: 10.5281/zenodo.4471118.
- [25] T. Srinivasarao, "Iris Flower Classification Using Machine Learning," *International Journal of All Research Education and Scientific Methods (IJARESM)*, vol. 9, no. 6, pp. 2–10, Mar. 2022.
- [26] A. Y. Muaad, A. Y. Abdullah, J. Davanagere, Hanumanthappa, J. V. Benifa, A. Alabrah, N. S. Amerah, M. Ahmed, D. Pushpa, M. A. Al-antari, and T. M. Alfakih, "Artificial Intelligence-Based Approach for Misogyny and Sarcasm Detection from Arabic Texts," *Computational Intelligence and Neuroscience*, vol. 2022, no. 7937667, pp. 1–9, Mar. 2022, doi: 10.1155/2022/7937667.
- [27] M. Hort, Z. Chen, J. M. Zhang, M. Harman, and F. Sarro, "Bias Mitigation for Machine Learning Classifiers: A Comprehensive Survey," *ACM J. Responsib. Comput.*, vol. 1, no. 2, pp. 1–52, Jun. 2024, doi: 10.1145/3631326.
- [28] M. Medvedeva, M. Vols, and M. Wieling, "Using machine learning to predict decisions of the European Court of Human Rights," *Artificial Intelligence and Law*, vol. 28, no. 2, pp. 237–266, Jun. 2020, doi: 10.1007/s10506-019-09255-y.
- [29] Minjung Park and Sangmi Chai, "AI Model for Predicting Legal Judgments to Improve Accuracy and Explainability of Online Privacy Invasion Cases," *Applied Sciences*, vol. 11, no. 23, pp. 1–16, Nov. 2021, doi: 10.3390/app112311080.