

PREDICTIVE BASED ANALYSIS OF CRIMINAL DATA (FIRS)

^{*1}Muhammad Khalid, ²Dr. Kareem Ullah^{*1}*Department of Computer Sciences, Baba Guru Nanak University, Nankana Sahib, Pakistan*²*Department of Computer Sciences, University of Agriculture, Faisalabad, Pakistan*^{*1}mkhalidns@gmail.com ²kareemkwl@yahoo.comDOI: <https://doi.org/10.5281/zenodo.20596082>**Keywords**

Component; Data Mining; Classification; Clustering; Crime; Security; Naive Bayes

Article History

Received on 28 Nov, 2025

Accepted on 25 Dec 2025

Published on 26 Dec 2025

Copyright @Author

Corresponding Author:

Muhammad Khalid

Abstract

Uses of Information Communication Technology (ICT) have significantly changed the working of almost every field. Law enforcement authorities also use these technological developments to maintain the public security. Data regarding criminal activities is collected and stored in large quantity which needs to be analyzed for finding different patterns. This analysis unlocked many secret of stored data which are many times ignored. Different Data Mining (DM) techniques have been developed for analyzing these data to find spot trends. So, there is a need to implement data mining techniques on criminal data to identify the crime patterns which are significantly helpful for law enforcement agencies to plan and control the security. In this paper, a complete analysis is performed by using predictive analytics techniques to analyze the criminal data to discover the crime trends. A framework has been purposed for predictive based analysis which is able to process large amount of criminal data and identify the hidden crime pattern and trends. Data is collected from local police to extract the important attributes. Purposed model used clustering and classification techniques, k-mean, J48 Decision Tree and Naive Bayes (NB) to explore the hidden crime profiling. These explorations are used for future prediction of crime occurrences. Later these predictions will facilitate the law enforcement authorities while making the security policies and deployment of security personals.

I. Introduction

It is common phenomenon of the world that security of every citizen is must be ensured by the state and state organs. Enormous challenges have been faced by every country of the civilized world to maintain security of their citizens. Crime ratio has been increased significantly and authorities need much effort to combat crimes. One of important way is the proper analysis of criminal data to identify the crime patterns. Data Mining (DM) is said to be one of the effective technique for extracting the hidden relationship, pattern and trends from huge volume of data which are used for future prediction and help the institutions to make their decision making process more effective with higher confidence level [1]. For this purpose, data analytic techniques used significantly for profiling the criminals, crime mapping and detection for analysis of certain crime type with different patterns and to forecast of occurring crime in future.

Crime data analytics have some important challenges to develop or purposed such an efficient crime identification tool to measure and analyzing the committed crime with different attributes i.e. time, reason and place. The importance challenges of analyzing crimes includes the storing of large amount of criminal data, inconsistency of criminal data, very difficult to obtain from local police and accuracy of predictions depends on the data set [2].

There are different types of data mining approaches used for processing data to extract the hidden information and relationships. Association approach discovers the items that occurring frequently in the dataset and describes the patterns [3]. But this technique greatly applied in the detection of network intrusion from the user's history. Classification technique explores some common properties from different entities of crime and arranged them in classes which are pre-defined. It mainly used to discover the email spamming source. Clustering of data means partitioning data and makes these partitions into subclasses also called clusters. It also helps the users to

understand structure of data set. Clustering is not an algorithm but a way to solve the problem. It uses small group notions to form a cluster which is useful to discover different similar patterns of data [4]. Prediction approach is used to predict the future outcomes from the past data after analyzing different pattern, trends and relation [5]. This approach is used with the combination of different data mining approaches. Sequence technique is used to identify the homogeneity or the happening of certain events regularly. The acquired information suggests automatically that when a certain event will happen, its related event will must be happened. Decision Tree is mostly related to the other techniques of data mining like prediction, classification and association. DT primarily used to support and select required data from the whole data set [6].

II. Related Work

A framework for mining criminal data is presented by [7] to draw conclusion on the basis of analysis which have been done on Coplink project. This project contains 1.3 million records of criminal. Different gangs are identified involving in murders, assaults and distribution of drugs. Analysis of these criminal networks has been done to extract the relationships between criminal incidents. [8] presented the behaviour approach to study the criminals and other statistics. Clustering approach of data mining has been used to analyzing the criminal activities to get the get clear picture of criminals. This method enables the police force to identify the spots in which the behaviour of criminal has been changed. [9] described a novel idea of analyzing crime patterns at micro level. This idea used edges of crimes for analysis at micro level to differentiate crime in urban areas. Data of Columbian and Canadian municipality compared by constructing it in different zones and adjacent neighbourhood's areas of edges are identified. [10] stated the applicability of predictive analytics and data mining tools in the security institutions for categorizing the crime pattern and the behaviour of criminal's involved in different types of crimes. It discuss the main characteristics of these tools, is

their ability to handle large amount of data and to identify the crime trends to combat the crimes by the early actions of security agencies. [11] used the k-mean clustering techniques on the criminal data for identification the patterns in crimes and then validate the findings through machine learning. Semi-supervised method also used to increase the accuracy of extracted knowledge. A weighted scheme has also been developed for implementation of data mining with geo-spatial plot the crime for countering the terrorism and to maintain security. [12] described that traditional data mining approach discovered spatial trends from the criminal data. A dynamic pattern framework analysis (DPA) presented to identify related type of patterns like changes in spatial patterns in some specific time and similarity in some different duration. The purposed framework (DPA) used to forecast the most likely crimes to happen. [13] stated the efficient methods of data mining on crime related data. These methods found the criminal activities and identify the criminals of the bases of historical data. But there are some problem faced for implementing such methods as the crime data is not consistently collected. [14] proposed an integrated model called PerpSearch which required crime description as input. These descriptions include crime type, crime location and physical characteristics of criminals. For crime detection, proposed model process these parameters by using profiling, analysis of social network, patterns of crimes and some physical similarities. These processing determine spot of criminals and identification of criminals.[15] introduced graph based approach for mining datasets to identify correlation between different criminal activities. This approach discovers the knowledge through analysis of criminal dataset. Datasets firstly transformed into directed graphs, mostly influenced edges are extracted for performing correlational analysis. [16] discussed the theory of crime patterns and described that the patterns in the spatial movement of criminals are very important as criminal do crime in their own spatial areas. This concept is used for

future prediction within the spatial areas of criminals. K means based clustering was used to calculate the activities of offenders.[17] focused on the spatial distribution and statistical analysis of crime pattern. Spatial method of crime distribution is used for identifying the central tendencies and dispersion degree from criminal dataset. Then hierarchical clustering and time space analysis has been implemented for visualizing the crime clusters. [18] stated the limitations of techniques used for discovering the crime pattern is integration. The important limitation is the integration to distinguish the random and clustered crime patterns. The integrated approach used k-function in both random and clustered methods for analysis. [19] experimented different types of crime across multiple location in a province of South Africa. An effective mining methodology is used in quite simple way to recognize the crime ratio. This methodology uses batch-merge prototype which builds on Fp-Growth algorithm and called CitiSafe algorithm. [20](Babakura et al., 2014) presented the application of classification on the criminal data for predicting the crime for different states of USA. Criminal data must be real in its nature and must be analyzed properly. Two algorithms of classification Naive Bayesian and 2nd are back propagation is used predicting the theory of crime. [21] presented the use of computer science approach and criminal activities to develop mining model for identification of crime. This model not mainly emphasized on the occurrence of crime, offender's background or any other criminal intentions but its chief objective to consider the factors of occurring these crimes. [22] presented the CRIME TRACER approach for the analysis of spatial crime and make prediction. This approach uses probabilistic method to identify the offender's behaviour and their activities. Crime theory stated that the opportunity to do crime must be encountered to combat the crimes. [23] stated that crime analysis is very needful for law organization to combat the crime situation. Data of crime activities is must be collected and used for the future crime prediction. An

intelligent system for crime analysis developed to overcome the challenges faced by data analysts.

III. Research Methodology

According to data mining (MD) ranks as an interface tool between the statistical and other areas of knowledge. This method stands out for its ability to discover patterns with some practical meanings for user. The author also shows that you should stick to the fact that the misuse of this tool and non-compliance with their application criteria and analysis can lead to bias of results and, consequently,

to a wrong interpretation of the same [24]. This is an overall process with the aim of extracting information from large databases, using any procedure automatic (mathematical or computational) [25].

The KDD process requires direct dialogue with the knowledge previous researcher in decision making on the results found by the method. A "Figure 1" highlights the steps of KDD necessary to pattern discovery process on the bench data.

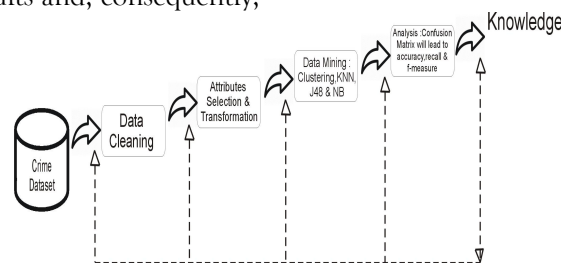


Figure 1: Steps of KDD

The data employed in this study was collected from the local police and prosecution department. Manual data is collected and computerized it as per the requirement of

this research. Table-1 enlists the attributes extracted from the manual data.

Table I: List of Attributes

Sr. No:	Name	Description
1	CaseID	Numeric
2	Case_Number	Numeric
3	Date	Date
4	Time	Time
5	Day	Monday, Tuesday, Wednesday, Thursday, Friday, Saturday, Sunday
6	Primary Type of crime	String
7	Description	String
8	Location Description	String
9	Arrest	True, False
10	Domestic	False, True

More than 6500 records have recorded; however, still there are a number of records stored manually awaiting to be entered to the automated system. Although theoretically large volume of data is more important to train data mining models, due to time constraint to pre-process the data, a sample data containing 2983 records are taken for this study. After the pre-processing stage, the

data were exported to CSV format (comma separated). A "Figure 2" illustrates the creation of the file where they were inserted the header Mining containing all categorical attributes previously announced. After standardized the .CSV file has been converted to file .Arff by using some conversion tool.

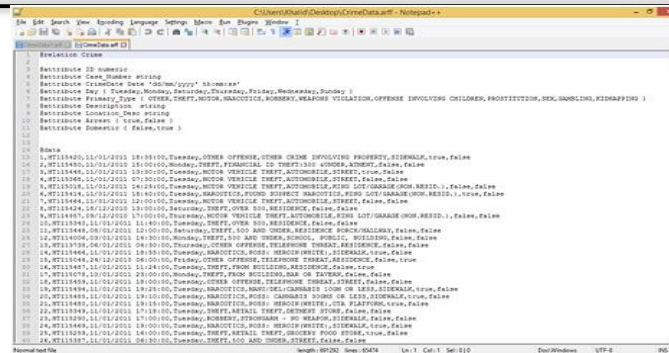


Figure 2: Creation of Arff file

Data Mining Tools

Currently there are different data mining and knowledge discovery tools available, such as Waikato Environment for Knowledge Analysis (WEKA), Rapid Miner, IBM SPSS Modeller and etc. These tools give different approaches and algorithms that used data in an efficient manner and provide accurate results. WEKA is considered to be a best tool to run the K-Mean, J48 and Naive Bayes (NB) algorithms. Weka is suitable for of pc learning algorithms for mining data. Many algorithms of data mining can be utilized instantly on any data set. This comprises all the tools for pre-processing, clustering, classification, association, regression and visualization of data [26].

This tool has a more comprehensive methodology for data preparation, data mining, and interpretation of results like k-means, DT J48 and NB. Below a brief discussion about the mining algorithm that are applied in this study.

Main Algorithm Used

The proposed model for the analysis includes clustering the data initially so as to group similar attributes together. How these algorithms works are explained below:

1. Clustering

The grouping task not implement algorithm supervised where the label of the class of each sample Training is not known, and the number or set of classes to be trained cannot be known a priori. At the grouping the objects are incorporated so that the similarity is maximum within each minimum cluster between instances of clusters differences [27]. According to [28] this feature aims to target a set of data in number of homogeneous subgroups or clustering. The fundamental difference between the formation of clusters and the classification is that there are classes in clustering default to sort the records under study. The records are built according to their similarities basic, that is, when it is desired to form groups, selects a

set of attributes (variables) and function the similarity of these characteristics groups are formed.

The criteria for selecting groups instead of other techniques such as classification and association, is because that other studies involving criminal data have already been developed using the cluster technique, generating quite interesting results. Another criterion is the fact of clustering technique to be unattended, i.e. unknown labels of classes of training data set and the main objective is establishing the existence of classes or groups of data.

• K-mean

It is a partition algorithm that search directly for great divide (or approximately optimal) of n elements without the need for hierarchical associations. Serve exclusively to group n Records in k clusters (Burke and Kayem, 2014). Proposed in 1967 by J. Macqueen, this is one of the algorithms most known and used, as well as being the one with the most number of variations [4]. The algorithm starts with the choice of K elements that formed the initial seeds. This choice it can be done in many ways, including:

- Selecting k premieres observations
- Selecting k observations randomly
- Choosing so k observations their values are quite different.

For example, when grouping a population into three groups according to the height individuals could choose a low individual height, a medium height and high.

The criterion of choice of k-mean occurs from other jobs involving mining crime data already use this algorithm to generate clusters. Another significant criterion is efficiency the method of treating large volumes of data, in addition to its computational simplicity, being easy algorithm application.

- **Euclidean distance**

Among the Euclidean distance measurement is the most known. To calculate the distance between two objects (i and j) with dimension p (number of variable), it is defined by:

Equation 1

$$d(i, j) = \sqrt{\sum_{l=1}^p (x_{il} - x_{jl})^2}$$

As a widely used derivation is Distance Euclidean Medium (DEM), where the sum of the difference square by the number element involved:

Equation 2

$$d(i, j) = \sqrt{\frac{\sum_{l=1}^p (x_{il} - x_{jl})^2}{p}}$$

Classification

Classification consists of predicting a certain outcome based on a given input. In order to predict the outcome, the algorithm processes a training set containing a set of attributes and the respective outcome, usually called goal or prediction attribute (Voznika and Viana n.d.). The algorithm tries to discover relationships between the attributes that would make it possible to predict the outcome.

- **Decision tree J48 algorithm**

The algorithm J48 classifier is a simple selection tree for classification. Decision trees use the heuristic called recursive partitioning. It uses the approach of divide and conquer where it uses the feature values to split into smaller subsets of similar classes. The whole dataset is represented by a node and by using the most predictive target class the data is divided into distinct classes which will be the tree branches. The divide and conquer

approach is used until a stopping criterion is reached. For splitting the data C5.0 uses entropy for measuring of purity. The algorithm uses a training data set initially to train the data and new test data is classified after that (Yamuna and Bhuvaneshwari, 2012).

- **Naïve Bayes classification algorithm**

Naive Bayes uses data about prior events to estimate the probability of future events. For each and every class label, multiply all the probabilistic conditions together for each feature, specified in the label of class. To implement the classifier, calculate all the conditional probabilities individually for each class label, of every feature, $p(F_i|C_j)$ and multiply at the side of the prior

probability for the label $p(C_j)$. The label of any high product, we got is actually the label which the classifier is returned (Tolan and Soliman, 2015). Equation-3 is the equation of Naive Bayes.

Equation 3

$$\text{classify}(f_1, \dots, f_n) = \arg \max p(C=c) \prod_{i=1}^n p(F_i=f_i|C=c)$$

Models testing and validation

Weka has the facility to extract and test the classifier accuracy on disjoint cases randomly. It provides three options to partition the dataset into training and test data. These are preparing files for training and testing data set, cross validation and percentage split (Al-rajab and Lu, 2016).

To test the models explained above, data set split in percentage 80 training set and 20 test set. The data set was randomised first. The supervised learning has two processes that is learning (training) the model by using training data and then testing the model using unseen data to assess for accuracy. The process is shown in Figure-3 below:

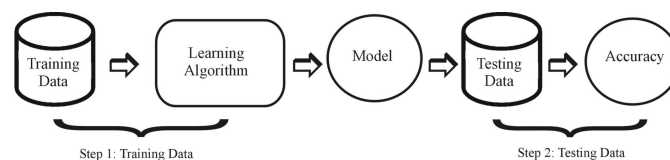


Figure 3: Training and testing data process

The purpose of this study was to explore the applicability of data mining techniques in the process of crime prevention. In doing so, 2983 sample data were taken randomly from criminal record database of the Local Police to build and test the model.

Result and Discussion

In this study, the prepared data set has been put into weka analyzer and got the following results. Table-2 shows the happening of crime in different weekdays.

Table II: *Result Of Day Wise Happening Of Crimes*

Sr. No:	Name	Description
1	Monday	346
2	Tuesday	350
3	Wednesday	513
4	Thursday	440
5	Friday	463
6	Saturday	411
7	Sunday	370

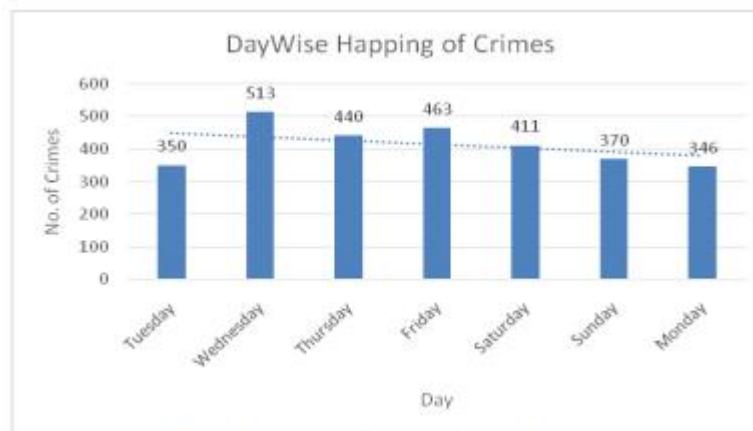
**Chart 1:** *Result of day wise happening of crimes*

Table-2 shows the result of happening of crime day wise where as the chart-1 shows graphical view of day wise crime situation. The result interestingly revealed the fact that Wednesday, Friday and Thursday are the day on which crime rate is very high and authorities should be more careful in that day.

Table III: *Result of day wise happening of crimes*

Sr. No:	Crime Type	Number of Occurrences	%
1	Theft	773	25.9
2	Narcotics	1383	46.4
3	Robbery	456	15.3
4	Weapon violation	88	3.0
5	Child offence	82	2.7
6	Prostitution	93	3.1
7	Gambling	5	0.2
	Kidnapping	13	0.4
Total		2893	

For all types of crimes simple k-mean is used together with the Euclidean distance function. Table-3 illustrates the happening of different types of crimes which shows percentage of theft (25.9%), Narcotics (46.4%), Robbery (15.3%), Weapon violation (3.0 %), Child offence (2.7 %), Prostitution (3.1 %), Gambling (0.2 %) and Kidnapping (0.4 %) of crimes.

Table IV: *Clustering using k-mean algorithm*

Number of iterations: 4
 Within cluster sum of squared errors: 14713.861723109447
 Missing values globally replaced with mean/mode

Cluster metadata		Cluster				
Instance	Full Data (1043393)	0 (1043393)	1 (335)	2 (754)	3 (395)	4 (563)
ID	1447	1446, 1448	015, 0407	1480, 4093	1476, 0407	1770, 7438
Class Number	01114420	01114420	01114420	01114420	01114420	01114701
Date (D-M-Y)	14/05/2013	14/05/2013	14/05/2013	22/05/2013	14/05/2013	17/05/2013
Time	12:00:00 PM	12:00:00 PM	12:00:00 PM	7:00:00 PM	4:00:00 PM	9:00:00 AM
Day	Saturday	Saturday	Friday	Saturday	Thursday	Monday
Primary Type	WAGCOTICE	WAGCOTICE	OTHER OFFENSE	WAGCOTICE	OTHER OFFENSE	OTHER OFFENSE
Description	POST: CURBAGEY DONE ON LANE	POST: CURBAGEY DONE ON LANE	TELEPHONE TALKING	POST: CURBAGEY DONE ON LANE	TELEPHONE TALKING	TELEPHONE TALKING
Location Description	STREET	STREET	STREET	STREET	APARTMENT	REFERENCE
Address	None	None	None	None	None	None
Domestic	False	False	False	False	False	False

Time taken to build model (full training data) : 34.11 seconds

=== Model and evaluation on training set ===

Clustered Instances

0	1043393 (100%)
1	335 (0%)
2	754 (0%)
3	395 (0%)
4	563 (0%)

Table-4 illustrates the result of applying clustering algorithm k-mean which shows that (100%) clustering is done and the task of clustering is completed successfully.

Results of J48 algorithm

Decision tree is very powerful and widely used data mining method for classification and prediction [34].

Table V: Results of J48 decision tree

=== Classifier model (full training set) ===

Time taken to build model: 1.2 seconds

=== Stratified cross-validation ===

=== Summary ===

Correctly Classified Instances	2666	92.1535 %
Incorrectly Classified Instances	227	7.8465 %
Kappa statistic	0.4246	
Mean absolute error	0.109	
Root mean squared error	0.2343	
Relative absolute error	60.6912 %	
Root relative squared error	78.2642 %	
Total Number of Instances	2893	
Ignored Class Unknown Instances	1042547	

=== Detailed Accuracy By Class ===

TP Rate	FP Rate	Precision	Recall	F-Measure	ROC Area	Class
0.986	0.663	0.931	0.986	0.958	0.84	false
0.337	0.014	0.729	0.337	0.461	0.958	true
Weighted Avg.	0.922	0.599	0.911	0.922	0.908	0.851

After applying this algorithm the results are show in Table-5. The correctly classified instances are 92.1535% and the kappa statistic is 0.4246 which is near one and these figures are a good indicator of a good model.

Table VI: Confusion Matrix

	a	b	← classified as
2569	36		a = false
191	97		b = true

A confusion matrix illustrates the accuracy of the strategy of classification concern. This matrix presented the actual and predicted classification completed by the classification model [35]. It is also used to evaluate the performance of such methods and the Table-6 shows the

matrix for this information. Diagonal elements of above matrix $2569+97=2666$ shows the instances correctly classified while other $36+191=227$ shows the instances not classified correctly. Figure-4 visualized the decision tree for J48 classifier algorithm by using weka.

=== Classifier model (full training set) ===

Time taken to build model: 0.86 seconds

=== Stratified cross-validation ===

=== Summary ===

Correctly Classified Instances	2468	85.3094 %
Incorrectly Classified Instances	425	14.6906 %
Kappa statistic	0.4801	
Mean absolute error	0.1491	
Root mean squared error	0.3533	
Relative absolute error	83.0409 %	
Root relative squared error	118.0024 %	
Total Number of Instances	2893	
Ignored Class Unknown Instances	1042547	

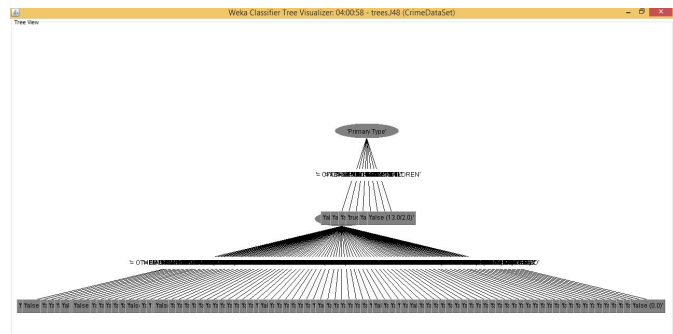


Figure 4: Classifier J48 Tree

Results of Naive Bayes algorithm

The Naive Bayes is straightforward probabilistic classifier algorithms that compute the set of probabilities with the aid of counting the frequency and mixtures of values in a given knowledge set. The algorithm is based on Bayes theorem and all the attributes are assumed as

independent given the value to class variable. After applying the Naive Bayes algorithm the results are displayed in Table-7. The correctly labelled instances are (85.30%) which is an efficient indicator and the kappa is 0.48 which is also a good indicator for purposed model.

Table VII: Confusion Matrix for Naive Bayes

.=== Confusion Matrix ===

a	b	←classified as
2206	399	a = false
26	262	b = true

Tabl-8 illustrates the confusion matrix of Naive Bayes algorithm which shows the information in diagonal elements of above matrix $2206+262=2468$ shows the instances correctly classified while other $399+26=425$ shows the instances not classified correctly

I. Conclusion and future work

This paper presented a purposed model for the knowledge discovery patterns from criminal data. There is a growing breakthrough in crime rates and violence which causes great concern by all law enforcement agencies. So the analysis of this criminal data now becomes an important tool for prevention of possible

occurrence of crimes. Data mining techniques clustering and classification have been used to understand the happening of crimes of some frequent form of day, or every month. This analysis make it easier to create a diagnosis of crime to the planning of new public policy fighting crime, specifically for each district, through lectures, educational blitz, creating units pacifying police, overt rounds, social programs, among many other possibilities. As future work, it is intended to include new months and years in the database for new algorithms to be used and can provide results higher than those found.

References

- [1] R. T. Scarfe and R. J. Shortland, "Data mining applications in BT," *Knowledge Discovery in Databases, [IEE Colloquium on]*, vol. 1995. IEE, p. 5/1-5/4, 1995.
- [2] Z. S. Zubi and A. A. Mahmmud, "Using Data Mining Techniques to Analyze Crime patterns in the Libyan National Crime," pp. 79-85, 1999.
- [3] S. S. Khan, M. A. Peer, and S. M. K. Quadri, "Comparative study of streaming data mining techniques," *2014 Int. Conf. Comput. Sustain. Glob. Dev.*, pp. 209-214, 2014.
- [4] M. Sharma, "Z - CRIME: A data mining tool for the detection of suspicious criminal activities based on decision tree," in *2014 International Conference on Data Mining and Intelligent Computing, ICDMIC 2014*, 2014, pp. 1-6.
- [5] A. Abbasi, "Predicting Behavior," no. june, 2015.
- [6] S. Thammaboosadee, B. Watanapa, and N. Charoenkitkarn, "A framework of multi-stage classifier for identifying criminal law sentences," *Procedia Comput. Sci.*, vol. 13, pp. 56-59, 2012.
- [7] H. Chen, W. Chung, J. J. Xu, G. Wang, Y. Qin, and M. Chau, "Crime data mining: A general framework and some examples," *Computer (Long. Beach. Calif.)*, vol. 37, no. 4, pp. 50-56, 2004.
- [8] J. S. De Bruin, T. K. Cocx, W. a. Kusters, J. F. J. Laros, and J. N. Kok, "Data mining approaches to criminal career analysis," *Proc. - IEEE Int. Conf. Data Mining, ICDM*, pp. 171-177, 2006.
- [9] S. Gorham, "The Edge Effect," *Ecotone*, vol. 2, pp. 99-111, 2006.
- [10] C. McCue, "Data mining and predictive analytics in public safety and security," *IT Professional*, vol. 8, no. 4, pp. 12-18, Jul-2006.
- [11] S. Nath, "Crime pattern detection using data mining," *Web Intell. Intell. Agent Technol.*, vol. 1, no. 954, pp. 1-4, 2007.
- [12] K. Leong, J. Li, S. Chan, and V. Ng, "Dynamic Pattern Analysis Framework for cooperative crime prevention," *Comput. Support. Coop. Work Des. 2008. CSCWD 2008. 12th Int. Conf.*, pp. 1053-1058, 2008.
- [13] P. Thongtae and S. Srisuk, "An analysis of data mining applications in crime domain," *Proceedings - 8th IEEE International Conference on Computer and Information Technology Workshops, CIT Workshops 2008*, pp. 122-126, 2008.
- [14] L. Ding, D. Steil, M. Hudnall, B. Dixon, R. Smith, D. Brown, and A. Parrish, "PerpSearch: An integrated crime detection system," *2009 IEEE Int. Conf. Intell. Secur. Informatics*, pp. 161-163, 2009.
- [15] P. Phillips and I. Lee, "Mining top-k and bottom-k correlative crime patterns through graph representations," *Intell. Secur. Informatics, 2009. ISI ...*, pp. 25-30, 2009.
- [16] R. Frank, M. A. Andresen, C. Cheng, and P. Brantingham, "Finding Criminal Attractors Based on Offenders' Directionality of Crimes," *2011 Eur. Intell. Secur. Informatics Conf.*, pp. 86-93, 2011.
- [17] J.-H. Wang and C.-L. Lin, "An Association Model Based on Modus Operandi Mining for Implicit Crime Link Construction," in *2011 International Conference on Advances in Social Networks Analysis and Mining*, 2011, pp. 548-550.
- [18] L. Li, W. Dong, C. Tian, and W. Sun, "Point pattern analysis utilizing controlled randomization for police tactical planning," in *Proceedings of 2012 IEEE International Conference on Service Operations and Logistics, and Informatics, SOLI 2012*, 2012, pp. 13-18.
- [19] O. E. Isafiade and A. B. Bagula, "CitiSafe: Adaptive Spatial Pattern Knowledge Using Fp-Growth Algorithm for Crime Situation Recognition," *2013 IEEE 10th Int. Conf. Ubiquitous Intell. Comput. 2013 IEEE 10th Int. Conf. Auton. Trust. Comput.*, pp. 551-556, 2013.
- [20] A. Babakura, N. Sulaiman, and M. A. Yusuf, "Improved Method of Classification Algorithms for Crime Prediction," pp. 250-255, 2014.
- [21] S. Sathyadevan, M. S. Devan, and S. Surya Gangadharan, "Crime analysis and prediction using data mining," *Networks & Soft Computing (ICNSC), 2014 First International Conference on*, pp. 406-412, 2014.

- [22] M. A. Tayebi, M. Ester, U. Glasser, and P. L. Brantingham, "CRIMETRACER: Activity Space Based Crime Location Prediction," *Asonam*, no. Asonam, pp. 472-480, 2014.
- [23] I. Jayaweera, C. Sajeewa, T. Wijewardane, and I. Perera, "Crime Analytics: Analysis of Crimes Through Newspaper Articles," 2015.
- [24] R. Article, "Data mining: a literature review," vol. 22, no. 5, pp. 686-690, 2009.
- [25] M. C. Holmes and D. D. Comstock-davidson, "Data Mining and Expert Systems in Law Enforcement Agencies," *Inf. Syst.*, vol. VIII, no. 2, pp. 329-335, 2007.
- [26] W. Chunyu, W. Xuehua, and Z. Xujuan, "Research on the improved frequent predicate algorithm in the data mining of criminal cases," in *Proceedings of the 2008 IEEE International Conference on Information and Automation, ICIA 2008*, 2008, pp. 1531-1535.
- [27] A. Takizawa, "Classification and feature extraction of criminal occurrence points using CAEP with transductive clustering," *Procedia - Soc. Behav. Sci.*, vol. 21, pp. 83-92, 2011.
- [28] E. Ekinici and H. Takci, "Using authorship analysis techniques in forensic analysis of electronic mails," in *2012 20th Signal Processing and Communications Applications Conference (SIU)*, 2012, pp. 1-4.
- [29] M.-J. Burke and A. V. D. M. Kayem, "K-Anonymity for Privacy Preserving Crime Data Publishing in Resource Constrained Environments," in *2014 28th International Conference on Advanced Information Networking and Applications Workshops*, 2014, pp. 833-840.
- [30] F. Voznika and L. Viana, "Data mining classification," pp. 1-6.
- [31] S. Yamuna and N. S. Bhuvaneswari, "Datamining Techniques to Analyze and Predict Crimes," *Int. J. Eng. Sci.*, vol. 1, no. 2, pp. 243-247, 2012.
- [32] G. M. Tolan and O. S. Soliman, "An Experimental Study of Classification Algorithms for Terrorism Prediction," *Int. J. Knowl. Eng.*, vol. 1, no. 2, pp. 107-112, 2015.
- [33] M. M. Al-rajab and J. Lu, "A study on the most common algorithms implemented for cancer gene search and classifications," vol. 14, no. 2, pp. 159-176, 2016.
- [34] G. S. Linoff and M. J. A. Berry, *Data Mining Techniques: For Marketing, Sales, and Customer Support*. New York, NY, USA: John Wiley & Sons, Inc., 1997.
- [35] N. N. Sakhare, "Classification of Criminal data using J48 Algorithm Classification of Criminal Data Using J48-Decision Tree Algorithm," vol. 4, no. February, pp. 167-171, 2016.