

A Short Review on Prediction and Recommendation Techniques

Nivedita Adhik Patil¹ and S. U. Mane²

^{1,2}Department of Computer Science and Engineering, Rajarambapu Institute of Technology, Sakharale, MS-415414, India.
Shivaji University, Kolhapur

Email: patilnewa@gmail.com, Sandip.mane@ritindia.edu

Abstract— Machine learning and deep learning techniques are widely used in various domains to solve different problems. The prediction and recommendation type of problems exists in various domains. Due to the importance of such types of problems, obtaining solutions to such problems becomes a necessity. Researchers have used existing as well as developed various approaches to solving prediction and recommendation types of problems. This study presents a short review of various machine and deep-learning techniques to address the prediction and recommendation problems. The machine and deep learning techniques are reviewed and future research scope is discussed in this paper.

Index Terms— Machine Learning, Predictive Models, Artificial Intelligence, Deep Learning, SVM, Linear Regression, K-nearest neighbor, Naïve Bayes, Decision Tree, Random Forest, Xg-Booster, CNN, RNN.

I. INTRODUCTION

Worldwide, suggestion administrations have become significant because of the way that they support web-based business applications and different examination networks. Recommender frameworks have an enormous number of uses in many fields, including financial, schooling, and logical examination. Different experimental examinations have shown that recommender frameworks are more compelling and dependable than catchphrase-based web crawlers for removing helpful information from enormous measures of information. The issue of suggesting comparative logical articles in an academic local area is called a logical paper proposal. A logical paper proposal means to suggest new articles or traditional articles that match specialists' inclinations. It has turned into an appealing area of study since the quantity of academic papers increments dramatically. In the realm of PDAs, computerized advancements are changing and getting progressed step by step. The internet-based information has expanded tremendously.

In the world of smartphones, digital technologies are changing and getting advanced day by day. The online data has increased enormously. Machine Learning (ML) is a part of Artificial Intelligence (AI) this system provides the ability to learn from previous experience. In today's world of technology, almost everything is dependent on Artificial Intelligence and Machine Learning. Machine learning is used in many applications from healthcare to finances to E-Commerce to weather. Whereas predictive analysis is like predicting the future. Prediction and Recommendation systems are the subclass of machine learning, which generally work as rating or ranking products. The recommendation system helps users to get relevant results from the huge database. Using an

information filtering recommendation system provides personalized item recommendations to users. Recommended items help users in numerous decision-making tasks, which movie to watch, which location to visit, and how weather is likewise. Prediction involves using historical data to make predictions about future events or outcomes.

Machine learning models are trained on existing data to learn patterns and relationships, and then used to predict unknown or future data points. Recommendation systems are used to suggest items or content to users based on their preferences, historical behaviour, or similarities to other users. Recommendation algorithms analyse user data to provide personalized recommendations. There are ample examples of prediction and recommendation. Disease prediction, Multiclass prediction model for students, Criminal prediction, Tourism recommendation and Movie recommendation using machine learning etc.

The recommendation system splits into three parts based on its computation method, collaborative filtering, content-based filtering, and hybrid filtering. Collaborative filtering uses the collaborative ratings given by users to make recommendations. The content-based filtering method uses descriptive attributes of items to make recommendations. And the hybrid filtering method combines multiple filtering methods to get a list of items according to user preferences.

Machine learning models can analyse historical stock market data to predict future stock prices or market trends. These predictions help investors make informed decisions. We can use ML for sales forecasting. By analysing past sales data along with other relevant factors such as seasonality, marketing campaigns, and economic indicators, machine learning models can predict future sales volumes, allowing businesses to plan inventory, resources, and marketing strategies accordingly. Weather prediction can be done by using machine learning algorithms. Algorithms can analyse historical weather data, such as temperature, humidity, wind speed, and precipitation, to forecast future weather conditions. This information is crucial for planning activities like agriculture, transportation, and disaster management.

Online marketplaces use recommendation systems to suggest products to users based on their browsing history, purchase history, and preferences. These recommendations can enhance the user experience, increase sales, and improve customer satisfaction. Platforms like Netflix, Amazon Prime, and YouTube utilize recommendation systems to suggest movies, TV shows, or videos to users based on their viewing history, ratings, and preferences. This helps users discover new content and keeps them engaged on the platform. News websites and content aggregators use recommendation systems to suggest articles or news stories to users based on their reading history and interests. This enables users to discover relevant content and enhances engagement.

II. TECHNIQUES FOR PREDICTION AND RECOMMENDATION

A. Machine Learning Algorithms

- *Support Vector Machine*

It is a machine learning algorithm for regression and classification. It solves regression problems by utilizing both theoretical and numerical functions. However, for classification, the support vector machine is the most appropriate. Finding the best hyperplane in an N-dimensional space to help distinguish data points in different classes within the feature space is the aim of the Support Vector Machine (SVM) algorithm. High accuracy and large datasets are also supported by SVM. The number of features that are used determines the hyperplane's dimension. SVM also builds 3D and 2D models. A set of labelled training data, with each data point represented by a feature vector and its matching class label, is fed into the SVM algorithm. The hyperplane that maximizes the margin between the two classes is the one that SVM looks for. The margin can be defined as the separation between the nearest data points for each class (also known as support vectors) and the hyperplane.

The hyperplane that maximizes this margin is the ideal one. When data cannot be separated linearly, SVM can map the data into a higher-dimensional space where linear separation is feasible using a method known as the kernel trick. Radial basis function (RBF) kernels and polynomial kernels are examples of common kernel functions. New, unseen data points can be categorized according to which side of the hyperplane they fall on once the hyperplane has been established. SVM can be computationally costly, particularly when dealing with big datasets. Additionally, the regularization parameter and the kernel selection have an impact. SVMs have been successfully applied in various domains, including image classification, text categorization, bioinformatics, and finance, among others.

- *Logistic Regression*

Predicting the probability is the main application of logistic regression in classification problems. Because it uses the sigmoid function for the output for a given class and the linear regression function as the input, this technique is known as logistic regression. Since it forecasts the categorical dependent variable's output, the result could be True or False, Yes or No, or 0 or 1. An "S" shaped logistic function is fitted by logistic regression. The algorithm's curve indicates binary outcomes, like if a customer would purchase a product. A labelled training dataset, in which each data point is represented by a feature vector and a binary class label, is the input used by logistic regression, along with other supervised learning algorithms. In order to determine the predicted probability, the algorithm first fits a linear regression model to the input features and then applies the logistic (sigmoid) function to the linear combination of features. The predicted probability is then transformed using a decision boundary (usually set at 0.5) into a binary class prediction. The instance is classified as class 1 if the predicted probability is greater than or equal to 0.5; otherwise, it is classified as class 0. A loss function known as logistic loss, also referred to as cross-entropy loss, is used in logistic regression to measure the discrepancy between the true class labels and the predicted probabilities during training. For many binary classification issues, including spam identification, medical diagnosis, and credit risk assessment, logistic regression is frequently utilized. It is interpretable, computationally efficient, and reasonably simple. On complicated, non-linear datasets, logistic regression might not function as well as other linear models. More sophisticated methods, such as multilayer neural networks or support vector machines, are frequently applied in these situations.

- *K- Nearest Neighbor*

The classification algorithm K-Nearest Neighbour has many uses in data mining, pattern recognition, and intrusion detection. It is predicated on the notion that labels or values of related data points tend to be similar. The KNN algorithm saves the entire training dataset as a reference during the training phase. Using a selected distance metric, such as Euclidean distance, KNN determines the distance between each training example and the input data point before making predictions. The K-means algorithm divides data into tiny groups. Data points precisely consist of at least one cluster that is most fit for the evaluation of a big dataset.

A labelled training dataset, consisting of individual data points represented by a feature vector and matching class label, is fed into a KNN. The algorithm determines the distance between each data point in the training set and the input instance using a distance metric, such as the Euclidean distance. The number of nearest neighbors to consider is represented by the parameter "k". This hyperparameter is set by the user. While a large value of k may smooth out decision boundaries, a small value may result in noisy predictions. Using the selected distance metric, KNN finds the k-nearest neighbors for a new input instance. Next, using the majority class of these k-nearest neighbors as a guide, it predicts the class of the new instance. If $k=1$, predictions are made using the single nearest neighbor's class. Although KNN is comparatively simple to comprehend and apply, because it calculates the distances between each data point in the training set, it can be computationally costly for large datasets. Furthermore, if the data contains many irrelevant features or a high number of features, it might not function well.

Naïve Bayes

The probability is determined by this classification algorithm using the Bayes theorem. Large datasets are worked with using this model. Each feature in this algorithm operates independently of the others. For text classification, the Naïve Bays algorithm is applied with a high-dimensional training dataset. When determining the probabilities of noisy data used as input, it provides the highest accuracy. Because every feature in a dataset is independent of every other feature, it is known as naïve. For instance, red, spherical fruit is identified as an apple if the fruit is identified based on its color, shape, and taste. As a result, each characteristic independently and independently of the others indicates that it is an apple. The simplest and fastest machine learning algorithm for classifying dataset predictions is Naïve Bayes. Naïve Bayes can be applied to both multi-class and binary classifications. Compared to the other algorithms, it does well in multi-class predictions. It is an analogy classifier that makes a comparison between a training tuple and a training dataset.

Decision Tree

Is an algorithm for supervised learning that is primarily used to solve classification issues. This classifier has a structure resembling a tree, with internal nodes standing in for features, branches for choices or rules, and leaf nodes for results. We pose various kinds of questions at every decision tree node as we build the tree. We will determine the information gain that corresponds to the question posed. When building the tree, information gain determines which feature to split on at each stage. The procedure will continue until all offspring nodes are pure or until the information gain equals zero, with the split with the largest information gain being selected as the initial split. Here, impure denotes that the data is a mixture of all classes, whereas pure indicates that all the data

is part of the same class. The CART (Classification and Regression Tree) Ensemble learning algorithm is also included in the decision tree.

Random Forest

Feature selection, regression, classification, and outlier detection are just a few of the many applications that make extensive use of Random Forest's strength and versatility. To get the best model performance, it is necessary to pay attention to the hyperparameters, which include the depth of each individual tree and the number of trees in the forest. Techniques like cross-validation can be used to fine-tune these hyperparameters.

A labelled dataset with features (input variables) and corresponding labels (output variables) for each data point is necessary for the algorithm to function. Using replacement sampling, Random Forest divides the training data into several subsets. We refer to each subset as a "bootstrap sample." Decision trees are trained individually using these samples. A random subset of the initial feature set is chosen for every tree in the forest. This random selection increases the diversity of the model and lowers the correlation between the trees. The bootstrap sample and the randomly chosen subset of features are used to build each tree in the Random Forest. The lack of pruning and deep growth of the trees may cause overfitting of the training set. Each tree "votes" for a class when making a prediction for a classification task; the class with the majority vote becomes the final prediction. The final output for regression tasks is calculated by averaging the predictions made by each tree. Compared to a single decision tree, Random Forest lowers the chance of overfitting by combining multiple decision trees. Random Forest is robust to outliers and noisy data because it relies on the majority vote or averaging from multiple trees.

- **Xg-Booster**

XGBoost is a broad execution of slope support. We should talk about certain highlights of XGBoost that make it so appealing. XGBoost offers regularization, which permits you to control overfitting by presenting L1/L2 punishments on the loads and inclinations of each tree. This element is not accessible in numerous different executions of angle helping. One more element of XGBoost is its capacity to deal with scanty informational indexes utilizing the weighted quantile sketch calculation. This calculation permits us to manage non-no passages in the element framework while holding similar computational intricacy as different calculations like stochastic slope plummet. XGBoost likewise has a block structure for equal learning. It makes it simple to increase on multicore machines or groups. It likewise utilizes store mindfulness, which lessens memory use while preparing models with enormous datasets. At long last, XGBoost offers out-of-centre registering abilities utilizing plate-based information structures rather than in-memory ones during the calculation stage.

B. Deep Learning

Profound learning is a part of AI which depends on counterfeit brain organizations. It is equipped for learning complex examples and connections inside information. In profound learning, we do not have to program everything unequivocally. It has become progressively famous as of late because of the advances in handling power and the accessibility of enormous datasets. Since it depends on fake brain organizations (ANNs) otherwise called profound brain organizations (DNNs). These brain networks are enlivened by the construction and capability of the human cerebrum's natural neurons, and they are intended to gain a lot of information.

The critical trait of Profound Learning is the utilization of profound brain organizations, which have numerous layers of interconnected hubs. These organizations can learn complex portrayals of information by finding various levelled examples and elements in the information. Profound Gaining calculations can naturally gain and improve from information without the requirement for manual component designing. Profound Learning has made critical progress in different fields, including picture acknowledgement, regular language handling, discourse acknowledgement, and suggestion frameworks. A portion of the well-known Profound Learning designs incorporates Convolutional Brain Organizations (CNNs), Intermittent Brain Organizations (RNNs), and Profound Conviction Organizations (DBNs).

- **CNN (Convolutional Neural Network)**

The input, catching, order, and yield blocks of the graph represent the CNN model for notice pictures. With three main handling sections—the setups and introduction, CNN Learning and Representation, and the get test this image has a lot of information. Two sources of information are used in the setups and Instate section: the case and the arrangements. First and foremost, the case shows how many layers there are, their names, how many channels there are on each convolution layer, the order technique, the size of the information pictures, and the percentage of convolution layers. To sum up, the model's design illustrates several boundaries, including learning rate, small-scale bunch, weight rot, and energy. The nLmF-CNN CNN Learning and Perception component displays picture

highlights across multiple layers, including the information, convolution, ReLU, pool, fc, and softmax layers.. As displayed in fig.1.

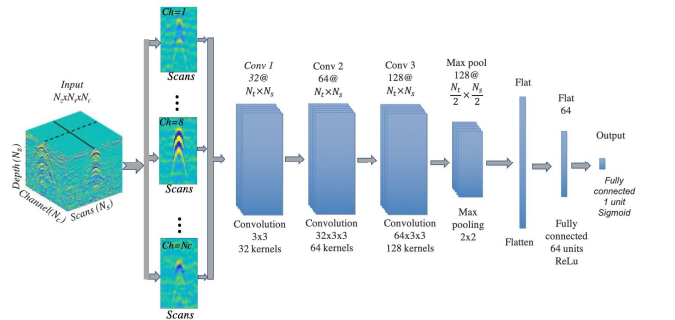


Fig1. CNN Architecture

It illustrates how learning status continues to have an impact. Convolution creates a result highlight map, also known as a convolved include (which may have a different size and profundity than the information highlight map), by removing tiles from the information include guide and applying channels to them to figure out new elements. Two boundaries are present in convolutions: The dimensions of the extracted tiles (typically 3x3 or 5x5 pixels). The depth of the highlight map result, which contrasts the quantity of channels used

The channels, which are networks of the same size as the tiles, move upward and across the data, including the framework of the guide, removing each related tile as they go. This process is known as a convolution. To introduce nonlinearity into the model after each convolution activity, the CNN modifies the convolved highlight using a Redressed Straight Unit (ReLU). When an x exceeds 0 and 0 when an x falls below 0, the ReLU capability, $F(x)=\max(0, x)$, yields x .

- **RNN**

Long Transient Memory Organization is a high-level RNN, a sequential organization that allows data to persist. It is prepared to handle the RNN-identified problem with vanishing inclination. RNN, or repeated brain network, is a technique used for careful memory. Let us say that when reading a book or viewing a film, you remember what happened in the previous scene. They collect historical data and use it to manage current information, much like RNNs do. One limitation of RNNs is that their vanishing slope makes it impossible for them to remember long-term scenarios. RNNs are specifically made to stay away from issues related to long-distance dependence. The first section decides if the information from the prior timestamp needs to be remembered or if it can be disregarded because it is not significant. In the following section, the cell attempts to forward new data from the contribution to this cell. In the third and last section, the cell transmits the updated data from the current timestamp to the subsequent timestamp. Entryways are these three elements that make up an RNN cell. The first part is called the "Result door," the second the "Information entryway," and the third the "Neglect door."

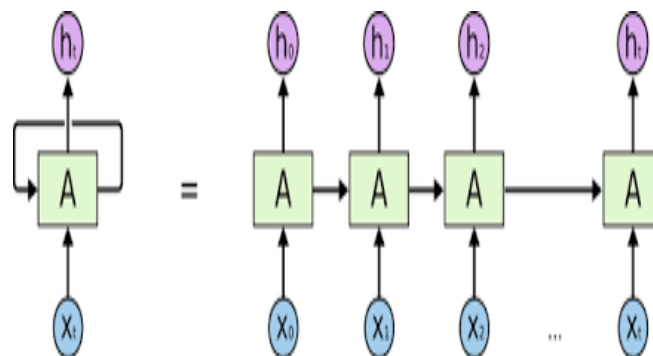


Fig.2: RNN Architecture

i. Input Gate

The objective of this step is to figure out what new data ought to be added to the organization's long-haul memory (cell state), given the past secret state and new info information.

ii. Forget Gate

Disregard door concludes the amount of past information will be neglected and the amount of the past information will be utilized in subsequent stages.

iii. Output Gate

The resulting door decides the worth of the following secret state.

III. APPLICATION OF MACHINE LEARNING AND DEEP LEARNING TECHNIQUES TO SOLVE DIFFERENT PREDICTION AND RECOMMENDATIONS

Using prediction and recommendation, we can find solutions to educational issues. [7] In order to increase the predictive accuracy and predict the final student grades for first semester courses, Siti Dianah Abdul Bujang studied and analyzed machine learning techniques. First, they made a comparison of six popular machine learning techniques' accuracy. After that, they created a multiclass prediction model using two feature selection techniques and oversampling Synthetic Minority Oversampling Technique (SMOTE) to reduce overfitting and misclassification results that came from imbalanced multi-classification. The highest f-measure of 99.5% is obtained in this case from the Random Forest results.

We can also use a machine learning algorithm to create a good example of a search engine. [8] According to Safia Kanwal, there is a vast amount of textual data on the cloud, including user reviews, news articles, eBooks, research articles, and personal blogs, but it can be challenging to locate the precise piece. The dataset, feature extraction methods, computational approaches, and evaluation metrics are the four primary facets of textual-based recommendation systems that are covered in this paper. Though there are other methods, they employed collaborative and content-based filtering in this instance.

By using machine learning algorithms, we can improve the yield of crops. In India, most people work as farmers. But they traditionally farm and take the same crop again and again, using the fertilizers in the wrong way, and because of this the yield gets affected, and damages the soil quality. [23] Nishchita K, states that using machine learning we can solve agricultural problems. Here for rainfall prediction, they used the Support Vector Machine algorithm and for crop prediction decision tree algorithm is used. This system will provide information about the required nutrients for the soil, required seeds for cultivation and the right crop for cultivation.

Healthcare is a major problem to tackle, and detecting the correct disease is very difficult using a machine learning algorithm it can be get predicted in the right way. Here [15] Noreen Fatima, this is a review paper in which different machine learning, deep learning, and data mining algorithms are reviewed.

There is abundant use of Machine learning in healthcare. We can also use it for alcohol disorder prediction. [12] Ali Ebrahimi explained it. This literature highlights five aspects of data collection sites, characteristics, type of dataset, data sampling and data pre-processing. The author stated that most studies use Support Vector Machine (SVM) as the main ML algorithm to predict Alcohol Use Disorder.

The daily rise in crimes indicates how important it is for people to take public security into account when they travel and relocate. [11] Wajiha Safat studied the growing rates of crime in Chicago and Los Angeles. In this work, a number of machine learning algorithms were employed, such as decision trees, logistic regression, random forests, logistic regression, support vector machines (SVM), k-nearest neighbors (KNN), eXtreme Gradient Boosting (XGBoost), and multilayer perceptrons (MLP). and found that these results, which are useful in directing law enforcement, allow for the quicker identification of criminal activity, hotspots with higher crime rates, and future trends with better predictive accuracy than with other techniques.

The stock market's extreme noise, complexity, and chaos make it a challenging maze to navigate. In order to forecast future stock prices or market trends, machine learning models can analyze historical stock market data. These forecasts assist investors in making wise choices.[13] Based on the outcomes of the training process, the Yaohu Lin states framework can recommend appropriate machine learning prediction techniques for every pattern. The investment plan is established using ensemble machine learning methodologies. China's stock market's observed results from 2000 to 2017 demonstrate that feature engineering produces accurate predictive results, with some trend patterns exhibiting prediction accuracy of over 60%.

Abbreviation: LR (Logistic Regression), SVM (Support Vector Machine), KNN (K- Nearest Neighbour), RF (Random Forest), CNN (Convolution Neural Network). NLP (Natural Language processing), LSTM (Long-Short Term Memory)

Here among all studied paper, we have mentioned some of them for analysis. For recommendation like students' grade prediction, we can use the SVM algorithm, which may go well. But for the critical data sets like stock

prediction or cancer prediction need to use the Random Forest GBDT or Neural Network algorithms. Here study says that one can find the solution for any kind of problem using recommendation and prediction of machine learning technique.

Reference	ML tech	Problem statement	Dataset used	Result & analysis
[7]	Decision Tree, Support Vector Machine, Naïve Bayes, K-Nearest Neighbour, Logistic Regression and Random Forest	Multiclass prediction model for student grade prediction using machine learning	Student dataset with classification into 5 classes	This study states that for student grade prediction Random Forest give the highest f-measure of 99.5%
[8]	NLP, CNN, Multi-Layer Perception.	A Review of Text-based Recommendation System	News, Movielens, Publications/ Journals or conferences, Advertisement, Query Log	Paper cover fundamental techniques used, available datasets evaluation metrics and algorithm approached used for text-based recommendation.
[23]	SVM, Decision Tree	Crop Prediction Using Machine Learning Approaches	i)Soil PH ii) Temperature iii) Humidity iv) Rainfall v) Crop data vi) NPK values	This study predicts crop, its essential NKP values and market price.
[15]	Machine Learning Techniques, Ensemble Techniques, Deep Learning Techniques	Prediction of Breast Cancer, Comparative Review of Machine Learning Techniques, and Their Analysis	Numeric and Image into MRI and Ultrasound	Paper states that SVM is the most suited ML algorithm in Prediction of breast cancer than other algorithms
[12]	SVM, CNN, Logistic regression, Radial support vector machine, Random Forest, Elastic net. Naïve bayes,	Prediction the Risk of Alcohol Use Disorder Using Machine Learning: A Systematic Literature Review	clinical reports and survey questionnaires, electronic health record	In this study several datasets with unique characteristics are found. Most of the studies use SVM for the prediction of Alcohol Use Disorder
[11]	LSTM, Exploratory data analysis, Forecasting with an ARIMA model	Empirical Analysis for Crime Prediction and Forecasting Using Machine Learning and Deep Learning Techniques	Chicago and Los Angeles datasets were used with time, location, type of crime with 22 attributes	In this study different ML algorithms were examined like LSTM and ARIMA modelling. This study achieved accuracy of 94% & 885 for Los Chicago and Angeles respectively.
[13]	Logistic regression, KNN, SVM, RF, GBDT (Gradient Boosted Decision Tree), LSTM	Stock trend Prediction Using Candlestick Charting and Ensemble Machine Learning Techniques with a Novelty Feature Engineering Scheme	Stock data	This study states that it automatically selects the prediction model for daily k-line patterns. RF and GBDT gives the more accuracy. SVM is not suited for such problems

IV. FUTURE SCOPE

Deep Learning for Recommendation: While traditional machine learning methods have shown promise in recommendation systems, deep learning techniques such as neural networks, especially deep collaborative filtering models, hold great potential to capture complex patterns and interactions in large-scale recommendation datasets. Future research could focus on developing more efficient and effective deep-learning models for recommendation tasks.

Explainable AI in Recommendation: Future work could explore methods to make recommendation models more interpretable and provide explanations for their recommendations, especially in domains where interpretability is crucial, such as healthcare or finance.

Reinforcement Learning for Recommendations: Reinforcement learning (RL) has shown success in other domains, such as gaming and robotics. Applying RL techniques to recommendation systems could enable more interactive and dynamic recommendations, allowing the system to adapt to changing user preferences over time.

Context-aware Recommendations: Incorporating contextual information such as user location, time of day, or user mood into recommendation models could lead to more relevant and personalized recommendations. Future work could explore techniques for seamlessly integrating contextual data into recommendation systems.

Cross-Domain Recommendations: Developing recommendation models that can effectively transfer knowledge across domains could be beneficial, especially when user preferences and item characteristics vary significantly between domains.

Multi-objective Recommendations: Most recommendation systems focus on optimizing a single objective, such as accuracy or click-through rate. Future research could explore multi-objective recommendation systems, considering diverse criteria such as novelty, diversity, and fairness.

Real-time Recommendations: As users increasingly demand real-time recommendations, future work could focus on developing efficient and low-latency recommendation models capable of handling large-scale, time-sensitive data.

Remember that the landscape of machine learning research is continually evolving, and new challenges and opportunities may arise beyond the scope of the points mentioned above. Researchers and practitioners in the field will continue to explore and innovate in prediction and recommendation using machine learning to create more effective and user-friendly systems.

V. CONCLUSION

Recommendation and prediction are the techniques which can provide you the solution for any kind of problem. This study presented a short review related to recommendation and prediction techniques from machine learning and deep learning. The study reviewed applications from various domains involving prediction and recommendation problems. Since 2019 to 2023 studies are considered for analysis. The possible future research scope is also presented based on the study.

REFERENCES

- [1] Liang Jiang, Lu Liu, Jingjing Yao⁴ and Leilei Shi “A hybrid recommendation model in social media based on deep emotion analysis and multi-source view fusion” *Journal of Cloud Computing: Advances, Systems and Applications*.
- [2] Yongfu Zha, Yongjian Zhang, Zhixin Liu and Yumin Dong “Self-Attention Based Time-Rating-Aware Context Recommender System” *Hindawi Computational Intelligence and Neuroscience* Volume 2022.
- [3] Chaohui Liu, Xianjin Kong, Xiang Li, and Tongxin Zhang “Collaborative Filtering Recommendation Algorithm Based on User Attributes and Item Score” *Hindawi Scientific Programming* Volume 2022.
- [4] Yue Wang, Zhaoxiang Qin, Jun Tang and Wei Zhang “Optimization of Digital Recommendation Service System for Tourist Attractions Based on Personalized Recommendation Algorithm” *Hindawi Journal of Function Spaces* Volume 2022, Article ID 1191419, 9 pages <https://doi.org/10.1155/2022/1191419>.
- [5] Jianjun Ni, Guangyi Tang, Tong Shen, Yu Cai and Weidong Cao “An Improved Sequential Recommendation Algorithm based on Short-Sequence Enhancement and Temporal Self-Attention Mechanism” *Hindawi Complexity* Volume 2022, Article ID 4275868, 15 pages <https://doi.org/10.1155/2022/4275868>.
- [6] Mariusz Kleć, Alicja Wiczorkowska, Krzysztof Szklanny and Włodzimierz Strus “Beyond the Big Five personality traits for music recommendation systems” Kleć et al. *EURASIP Journal on Audio, Speech, and Music Processing* (2023) 2023:4 <https://doi.org/10.1186/s13636-022-00269-0>.
- [7] Siti Dianah Abdul Bujang, Ali Selamat, Roliana Ibrahim, Enrique Herrera-Viedma, Hamido Fujita, Nor Azura Md. Ghani “Multiclass Prediction Model for Student Grade Prediction Using Machine Learning” *IEEE journal* published on July 2021 Digital Object Identifier 10.1109/ACCESS.2021.3093563
- [8] Safia Kanwal, Sidra Nawaz, Muhammad Kamran Malik, And Zubair Nawaz “A Review of Text-Based Recommendation Systems” *IEEE journal* published on march 2021 Digital Object Identifier 10.1109/ACCESS.2021.3059312
- [9] Mohammed A. Fouad, Wedad Hussein, Sherine Rady, Philip S. Yu, And Tarek F. Gharib “An Efficient Approach for Rational Next-Basket Recommendation ” *IEEE journal* published on 22 July 2022 Digital Object Identifier 10.1109/ACCESS.2022.3192396
- [10] Robert Kwieciński, Grzegorz Melniczak, And Tomasz Góreck “Comparison of Real-Time and Batch Job Recommendations” *IEEE journal* published on 3 march 2023 Digital Object Identifier 10.1109/ACCESS.2023.324935
- [11] Wajiha Safat, Sohail Asghar, And Saira Andleeb Gillani “Empirical Analysis for Crime Prediction and Forecasting Using Machine Learning and Deep Learning Techniques” *IEEE journal* published on May 2021 Digital Object Identifier 10.1109/ACCESS.2021.3078117
- [12] Ali Ebrahimi, Uffe Kock Wiil, Thomas Schmidt, Amin Naemi, Anette Søgaard Nielsen, Ghulam Mujtaba Shaikh, And Marjan Mansourvar “Predicting the Risk of Alcohol Use Disorder Using Machine Learning: A Systematic Literature Review” published on November 2021 Digital Object Identifier 10.1109/ACCESS.2021.3126777
- [13] Yaohu Lin, Shancun Liu, Haijun Yang, And Harris Wu “Stock Trend Prediction Using Candlestick Charting and Ensemble Machine Learning Techniques with a Novelty Feature Engineering Scheme” *IEEE journal* published on July 2021 Digital Object Identifier 10.1109/ACCESS.2021.3096825
- [14] Theofanis Kalampokas, Konstantinos Tziridis, Nikolaos Kalampokas, Alexandros Nikolaou, Eleni Vrochidou, And George A. Papakostas “A Holistic Approach on Airfare Price Prediction Using Machine Learning Techniques” *IEEE journal* published on May 2023 Digital Object Identifier 10.1109/ACCESS.2023.3274669

- [15] Noreen Fatima, Li Liu, Sha Hong, And Haroon Ahmed “Prediction of Breast Cancer, Comparative Review of Machine Learning Techniques, and Their Analysis” IEEE journal published on August 2020 Digital Object Identifier 10.1109/ACCESS.2020.3016715
- [16] Sofia Nudrat, Hikmat Ullah Khan, Saqib Iqbal, Mian Muhammad Talha, Fawaz Khaled Alarfaj, and Naif Almusallam “Users’ Rating Predictions Using Collaborating Filtering Based on Users and Items Similarity Measures” Hindawi Computational Intelligence and Neuroscience Volume 2022, Article ID 2347641
- [17] Chuang Zhang, Ming Wu, Jinyu Chen, Kaiyan Chen, Chi Zhang, Chao Xie, Bin Huang, And Zichen He “Weather Visibility Prediction Based on Multimodal Fusion” IEEE journal published on June 2019 Digital Object Identifier 10.1109/ACCESS.2019.2920865
- [18] Jibing Gong, Weixia Du, Huanhuan Li, Qing Li, Yi Zhao, Kailun Yang, And Ying Wang “Score Prediction Algorithm Combining Deep Learning and Matrix Factorization in Sensor Cloud Systems” IEEE journal published on April 2021 Digital Object Identifier 10.1109/ACCESS.2020.3035162
- [19] Triyanna Widiyaningtyas, Indriana Hidayah and Teguh B. Adji “User profile correlation based similarity (UPCSim) algorithm in movie recommendation system” journal of big data Widiyaningtyas et al. J Big Data (2021) 8:52 <https://doi.org/10.1186/s40537-021-00425-x>
- [20] Liang Jiang^{1,2}, Lu Liu^{3*}, Jingjing Yao⁴ and Leilei Shi “A hybrid recommendation model in social media based on deep emotion analysis and multi-source view fusion” Journal of Cloud Computing: Advances, Systems and Applications Jiang et al. Journal of Cloud Computing: Advances, Systems and Applications (2020) 9:57 <https://doi.org/10.1186/s13677-020-00199-2>
- [21] Xin Liu “Impact of Personalized Recommendation on Today’s News Communication through Algorithmic Mechanism in the New Media Era” By Hindwi journals published on July 2022 Hindawi Advances in Multimedia Volume 2022, Article ID 1284071, 8 pages <https://doi.org/10.1155/2022/1284071>
- [22] Zengchen Yu, Syed Umar Amin, Musaed Alhussein, And Zhihan Lv “Research on Disease Prediction Based on Improved DeepFM and IoMT” IEEE journal published on march 2021 Digital Object Identifier 10.1109/ACCESS.2021.3062687
- [23] Mahendra N, Dhanush Vishwakarma, Nischitha K, Ashwini, Manjuraju M. R “Crop Prediction using Machine Learning Approaches” Volume 09, Issue 08 (August 2020)