



Machine learning predictive model for an intelligent tourism recommendation system

E. Aldhahri¹ • N. Aljojo^{2*} • A. Tashkandi² • A. Alshutayri¹
B. Al-Subhi¹ • E. Al-Jedaani¹ • S. Al-Shmrani¹ • W. Al-Kaberi¹

¹Department of Computer Science and Artificial Intelligence, College of
Computer Science and Engineering, University of Jeddah

²College of Computer Science and Engineering, Department of Information
Systems and Technology, University of Jeddah, Jeddah, Saudi Arabia

Received 11 02 2024; accepted 01 15 2025

Available 10 31 2025

Abstract: Recommendation systems, powered by machine learning, are essential in offering personalized recommendations to consumers based on their preferences in different areas, including literature, educational programs, and lodging. In the field of recommendation systems, there are numerous techniques and strategies. Given that tourism plays a crucial role in driving the economies of regions and countries, there is an increasing desire to enhance the tourist experience by improving the way information is provided. Nevertheless, current research often fails to address the need for comprehensive manuals on activities and attractions while traveling. This project aims to fill this gap by developing a machine learning predictive model for an intelligent tourist recommendation system. The system is designed to help travelers choose the best routes for their travels. This study uses machine learning algorithms such as Naive Bayes, Decision Trees, and Linear Regression to analyze the “Tourism rating” dataset obtained from Kaggle. The dataset consists of 12 significant features. The results indicate that Linear Regression surpasses other methods, exhibiting greater predictive accuracy and decreased error rates. The importance of this study lies in its ability to offer customized suggestions and a wide range of

*Corresponding author.

E-mail address: nmaljojo@uj.edu.sa (N. Aljojo).

Peer Review under the responsibility of Universidad Nacional Autónoma de México.

choices to travelers, thereby improving their travel experiences by directing them towards the most suitable destinations and activities.

Keywords: Intelligent tourism, artificial intelligence, machine learning, decision trees.

1. Introduction

In recent years, there has been substantial growth in the Internet, the Internet of Things (IoT), Cloud computing, and Artificial Intelligence (AI) (Shi et al., 2020). Moreover, there has been a significant surge in the volume of material accessible on the internet. This data can be utilized to enhance individuals' lives by furnishing them with a greater array of choices. With the increasing demand for automation in all aspects of human life, recommendation systems are becoming increasingly valuable. Part of the reason for this is their evolutionary adaptation to better cater to the demands of their customers. A recommendation system is a type of computer software that analyzes a user's previous interactions with an item, the behavior of other users who also use the same item, and their interactions with each other to generate personalized recommendations (Kim & Lennon, 2012). Consequently, the field of information suggestion has experienced a tremendous increase in prominence. Many organizations, such as Netflix and Spotify, utilize recommender systems to ascertain the factors that attract users and to acquire a competitive edge (Seroussi, 2015). Moreover, recommendation systems can be found in various contexts, such as online retailers, where they suggest products to customers, and in systems that assist individuals in discovering books, training courses, accommodations, and trip destinations.

The recommendation systems methods can be classified into three categories: "Content-Based (CB)", "Collaborative Filtering (CF)", and "Hybrid Algorithms". The content-based recommendation system utilizes the user's profile information and the characteristics of the products they are interested in to generate recommendations. The collaborative filtering recommendation system utilizes a rating similarity matrix of neighboring users and products to construct user preferences. CF has been classified into memory-based, model-based, item-based, and user-based filtering. Memory-based filtering utilizes the complete rating matrix to generate preferences, whereas

model-based filtering divides the rating matrix into training and testing data. Moreover, user-based filtering relies on the similarities between users, whereas item-based filtering relies on the similarities between items.

A hybrid recommendation algorithm combines content-based (CB) and collaborative filtering (CF) methods. It utilizes both the user's knowledge and the rating matrix to build preferences and improve performance (Patel & Patel, 2020). There is no flawless system for making recommendations, as it is contingent upon the specific problem being addressed. Furthermore, CB may lack the ability to discern reliable information from both high-quality and low-quality sources (Umanets et al., 2014). Nevertheless, three potential issues can arise with collaborative filtering: the first-rater problem, sparsity, and the grey sheep dilemma. The concept of a "first-rater problem" refers to the situation where a newly introduced item cannot be recommended to users until at least one user has rated it. Sparsity refers to the condition in which the rating matrix contains a high proportion of vacant cells. The grey sheep problem refers to the phenomenon where consumers who hold uncommon viewpoints or preferences often do not receive accurate recommendations.

The applications for recommendation algorithms are vast. Zhou et al. (2020) utilized CB algorithms to examine the textual content of tourists' evaluations of beautiful places in Yunnan. They also incorporated seasonal and heat index data from various scenic spots in their analysis. Umanets et al. (2014) have utilized collaborative filtering (CF) to generate suggestions by considering both the user and the item. By using the information provided by users, the researchers effectively recommended a set of sites to other users with similar interests. The text is referenced by number 5. Nilashi et al. (2017) suggested an alternative usage of the CF recommendation algorithm, where researchers generated user photos using data provided by the users. The degree of attribution for each user thought was calculated using the fuzzy set approach, which was included in the ontology knowledge base during the development phase. The system processes a

user's recommendation request by searching the question bank for concept attribution values. It then picks and initializes the user's interest degree to complete the dynamic update of the user's picture.

Ramzan et al. (2019) employed three machine learning techniques —Naive Bayes, Bayesian networks, and support vector machines—to categorize users based on the demographic data provided. The tour operators had assembled a cohort of travelers who shared common interests. Unfortunately, relying solely on demographic data is generally inadequate for reliably predicting the rating in most situations. Additionally, tourism serves as a substantial source of income for numerous countries and regions. Presently, a 10% rise in GDP in a country has the capacity to create one job out of every eleven, through its direct, indirect, or induced impact on the tourism sector. According to estimates, the number of foreign visitors to the world is projected to reach 1800 million by 2030. This is an annual increase of 3.3 percent compared to the 25 million foreign visitors in 1950 [9]. Building upon the previous research, this study sought to enhance an intelligent system by incorporating machine learning algorithms. The objective was to optimize recommendations for tourism destinations while minimizing costs.

Apart from the Abstract and this current section, the remaining part of this paper is organized as follows: Section 3 presents related work, Section 4 outlines the research methodology, Section 5 discusses the findings, and Section 6 provides the conclusion.

2. Related work

Within the realm of Intelligent Tourism recommendation systems, a multitude of machine learning predictive models are currently under investigation (see Figure 1). Zhou et al. (2020) created a tourism recommendation algorithm that optimizes tourist sightseeing and tour route recommendations. This algorithm is based on text mining methods for constructing a travel plan for tourism purposes. It is considered to be one of the most important works in the field. The confirmation of tourists' interests and the tourist attraction feature attributes will be performed by a matching algorithm that utilizes clustering and mining algorithms to identify the most relevant matches. The schedule will contain information about the tourists' interests, the services provided by the tourism city and tourist attractions, the tour's schedule time, and the cost of the tour. Each of these details will be included in the schedule.

The recommendation algorithm also utilizes the McCulloch-Pitts (MP) nerve cell to generate the route chains for the tour. This is done to ensure the tour is as efficient as possible. To evaluate the performance and stability of the suggested algorithm, an experiment was conducted in the urban tourist attractions of Leshan city. The single-point analysis test and quantitative approach were employed to conduct the assessment. According to the findings, the proposed algorithm was able to deliver a higher level of satisfaction concerning the motive benefit while simultaneously demanding less time and space complexity.

A novel recommendation system for tourism was developed by Nilashi et al. (2017). This system utilizes clustering techniques, including Self-Organizing Map (SOM) and Expectation Maximization (EM), to group items. Finally, in order to improve the system's precision, they implemented the Hypergraph Partitioning Algorithm, also known as HGPA. ANFIS, which stands for Adaptive Neuro-Fuzzy Inference Systems, and SVR, which stands for Support Vector Regression, were employed to produce predictions. Additionally, Principal Component Analysis (PCA) was utilized by the system to minimize the number of dimensions (SVR) successfully.

An experimental study was conducted using the dataset provided by TripAdvisor. In order to assess the effectiveness of the clustering techniques, the research makes use of performance measurements such as Precision@7, F1, Coverage, and Mean Absolute Error. For the approach of recommendation predictions, the research makes use of the coefficient of determination (R^2) and the Mean Squared Error (MSE), both of which are metrics that demonstrate how well the two methods perform together. The findings demonstrated that the PCA-ANFIS-EM-HGPAP technique achieved the highest F1 value (0.922) while also yielding the lowest MAE (0.761) and MAPE (0.159) compared to the other proposed methods.

On the other hand, Ramzan et al. (2019) introduced a novel collaborative filtering (CF) recommendation strategy that is based on polarity detection. Additionally, opinion-based sentiment analysis is used to generate a hotel feature matrix. In order to acquire an understanding of how people feel about hotel features and how different sorts of guests are profiled (solo, family, couple, etc.), their approach integrates three analyses: semantic analysis, lexical analysis, and syntactic analysis. This enables them to gain a better understanding of how people perceive hotel features. Additionally, their system utilizes large datasets from the Hadoop platform and provides

recommendations for hotel classes based on the type of guest. This is accomplished through the utilization of fuzzy rules. They begin by analyzing the raw textual evaluations and organizing them before moving on to extracting features that are based on opinions. When they trained the system with ideal data, they achieved an accurate rate of 95 percent. Their experiment is based on a precision measure, which allows them to evaluate the proposed approach. This is even though the accuracy lowers to 74% when the dataset is huge.

In addition, Jeong et al. (2020) presented a system based on deep learning that utilizes machine learning classifiers to offer users various forms of tourism tailored to their personality. Three layers make up the system that has been presented. These layers include service provisioning, recommendation service, and adaptive definition. The recommendation service layer is responsible for generating suggestions for services based on the information the user fills out. In addition, the Myers-Briggs Type Indicator (MBTI) personality type is broken down into two distinct categories, namely introverts and extroverts, by this module. The system achieved roughly 83.33 percent for the result that it produced.

Two different recommendation models were developed by Mana and Sasipraba (2021) with the application of machine learning techniques. These models were a recommendation system for banking services and a recommendation system for movies. In the first model, the naive Bayesian classifier technique was utilized, and the Twitter dataset was used for both training and testing the recommendation system. Despite the movie recommendation system utilizing a hybrid recommendation model that incorporated both content-based and collaborative filtering recommendations, it was trained and evaluated using multiple datasets. The researchers discovered that these models are capable of providing appropriate recommendation models even when dealing with a substantial volume of data.

In their study, Cui et al. (2019) proposed a tourism recommendation system that is based on tag-based collaborative filtering algorithms. An experimental tourism website was employed as the source of the dataset, which contained information on users. According to the top-n recommendation concept, which recommends a large number of tourist attractions with the highest expected score for each tourist, the outcomes of the suggestion for one user ranged from 0.0150 percent to 0.0099 percent.

Abbasi-Moud et al. (2020) suggested a tourist recommendation system that provides individualized

recommendations to users based on the tastes of those users. The reviews left by users are a rich source of information that can be used to extract users' preferences. The comments are initially evaluated in order to identify the preferences of tourists. With the help of semantic clustering and sentiment analysis, the data was analyzed, and the features were retrieved from the aggregated user reviews of an attraction. After that, the proposed suggestion system will evaluate a user's preferences in relation to the characteristics of attractions to provide the user with the most relevant areas of interest. In addition, the algorithm utilizes contextual information, including time, place, and weather, to eliminate inappropriate items and enhance the quality of relevant ideas in the current context. After evaluating the system using data obtained from the TripAdvisor platform, the findings indicate that the suggested system outperforms the previously existing systems in terms of the F-measure criterion.

Al-Ghobari et al. (2021) proposed a Location-Aware Personalized Traveler Recommender System (LAPTA), which combines user preferences with GPS to provide tailored and location-aware recommendations. This enables the proposed system to offer recommendations that surpass those provided by typical recommender systems used in customer ratings. Names and classifications are assigned to the information that is gathered from Google locations by LAPTA. The proposed system utilizes the K-Nearest algorithm to pair the user's input with name and category tags, generating suggestions that are specifically tailored to the user's preferences. The system also makes suggestions based on surrounding notable places by utilizing Google's point of interest feature. This is all done with the intention of improving the system's usability. The trial findings demonstrated that LAPTA was capable of providing more reliable and accurate suggestions compared to the other recommendation applications investigated.

3. Research methodology

In this experiment, the technique of the Item-based Collaborative Filtering Engine recommender was utilized. The procedures that are carried out within the recommendation system are illustrated visually in Figure 1. It represents the outcome of user collaboration in selecting the order of places. When developing suggestions for new users, the recommendations used are derived from those of existing users. In order to accomplish this goal, collaborative filtering takes into account the opinions of users regarding a variety of locations and then makes

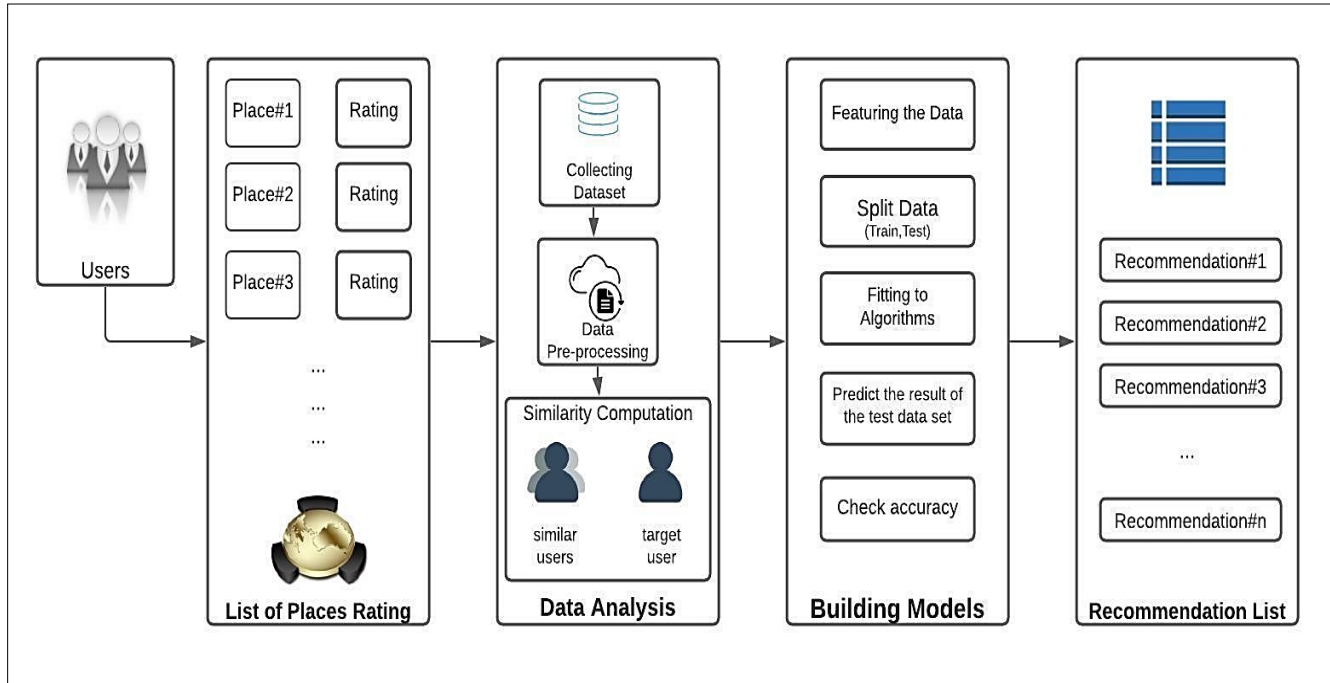


Figure 1

recommendations to each user regarding the perfect location based on their previous rating as well as the opinions of other users who are of a similar type (Parveez Saniya, 2020)

Item-based Nearest Neighbor: This process makes predictions based on similarities between items or places. A weighted total of the user's (U) ratings for items most similar to (I) is used to predict a user (U) and an item (I). Equation 1 shows the prediction equation (Parveez Saniya, 2020).

$$pred(u, i) = \frac{\sum_{j \in ratedItems(u)} sim(i, j) \bullet r_{uj}}{\sum_{j \in ratedItems(u)} sim(i, j)} \quad (1)$$

A recommendation system utilizing cosine similarity was developed. The cosine similarity metric is used to compare items/products regardless of their size. As a result, the cosine of an angle is calculated by measuring the distance between any two vectors in a multidimensional space. It is suitable for supporting papers of substantial size due to its dimensions. The cosine similarity equation is presented in equation 2 (Parveez Saniya, 2020).

$$Similarity(p, q) = \cos \theta = \frac{p \bullet q}{\|p\| \|q\|} = \frac{\sum_{i=1}^n p_i q_i}{\sqrt{\sum_{i=1}^n p_i^2} \sqrt{\sum_{i=1}^n q_i^2}} \quad (2)$$

3.1 Dataset and pre-processing

The dataset was obtained from Kaggle, consisting of two tables, each with more than 200 records. The following are the names of the two tables included in the dataset: tourism_rating.csv: This CSV file contains User_Id, Place_Id and Place_Rating.

tourism.csv: This CSV file contains Place_Id, Place_Name and other information about the place.

To begin, the researcher cleaned the data by replacing null values in the tourism rating.csv file with the mean of the column to which they belonged as part of the cleaning process.

The analysis of the dataset is performed through the following steps:

1. Finding the number of ratings in the training dataset to get the highest rating (1,2,3,4,5).
2. Compute the number of rated places per user and the rating number per place.
3. Check the Cold Start Problem: To find if there is any new user or place.
4. Compute Similarity Matrix: It is required to calculate the similarity between user profiles and places to create a recommendations list. The similarity between two data points is calculated using this type of matrix in Figure 3.

The feature extraction involves refining new features from the already existing features. The Features before extraction in the tourism_rating.csv file are: Uer_Id, Place_Id and rating. On the other hand, the features before extraction in the tourism.csv file are: Place_Id, Place_Name, Description, Category, City, Price, Rating, Time_Minutes, Coordinates, Latitude, Longitude, Unnamed: 11, and Unnamed: 12.

Here are the features after extraction:

- User_Id
- Place_Id
- gloabl_average
- similar_user_rating1
- similar_user_rating2
- similar_user_rating3
- similar_user_rating4
- similar_user_rating5
- similar_place_rating1
- similar_place_rating2
- similar_place_rating3
- similar_place_rating4
- similar_place_rating5
- user_average
- place_average
- rating

Three machine learning algorithms were used in the development of the proposed model for this research. They are addressed in detail in the following subsections:

Naive Bayes (NB) is a simple Machine Learning (ML) technique that is highly scalable, capable of efficiently handling expanding workloads. Additionally, NB is capable of handling enormous datasets due to the absence of iterative parameter estimation. Additionally, it is one of the simplest Bayesian network models. Simultaneously, NB can reach a great degree of accuracy. NB is based on Bayes' theorem, as seen in Equation 3, and predicts that (x) is a member of class (c) (Kolluri & Razia, 2020).

$$P(c|x) = \frac{P(x|c)P(c)}{P(x)}$$

$$P(c|X) = P(x_1|c) \times P(x_2|c) \times \dots \times P(x_n|c) \times P(c) \quad (3)$$

Where $P(c|x)$ is the posterior probability, $P(x|c)$ is the likelihood, $P(c)$ is the class prior probability, and $P(x)$ is the predictor prior probability.

The Decision Tree (DT) algorithm is a widely used classification technique. It is a tree structure made of if-then statements. Each internal node represents a condition, each leaf node represents a decision, and each

path represents a set of classification rules. Additionally, it is capable of managing a large volume of data. DT segments the tree using two matrices: Entropy and Gain. The entropy metric, depicted in Equation 4, quantifies the dataset's impure or random nature (Charbuty & Abdulazeez, 2021). The entropy value ranges from 0 to 1, with lower values being preferable. Additionally, as illustrated in Equation 5, the information Gain measure is the inverse of entropy, with a greater value being better (Charbuty & Abdulazeez, 2021). The Information Gain metric quantifies the amount of knowledge one has about the value of a random variable.

$$Entropy(S) = \sum_{i=1}^c P_i \log_2 P_i \quad (4)$$

$$Gain(S, A) = \sum_{v \in V(A)} \frac{|S_v|}{|S|} Entropy(S_v) \quad (5)$$

Linear Regression (LR) is another type of machine learning. It examines how the variables are linked together, categorizing them into "predictor" and "outcome" variables. It is also an extension of Pearson's correlation, which is what LR stands for. A simple LR does not have many features. In simple-LR, there is only one predictor variable that predicts one outcome variable, so there is only one outcome variable to look at. Multi-LR, on the other hand, utilizes a variety of predictor variables to predict multiple outcome variables. Equation 6 shows the multi-LR formula. Y_i is the outcome variable, and b_1 is the coefficient of x_{1i} , b_2 is the coefficient of the predictor variable of (x_{2i}) , b_n is the coefficient of the predictor variable of (x_{ni}) , and e_i is the error term (Olsen et al., 2020)

$$Y_i = (b_0 + b_{1x_{1i}} + b_{2x_{2i}} + \dots + b_{nx_{ni}}) + e_i \quad (6)$$

3.2 Model training and evaluation

According to the model's analysis, different partitions were utilized, with the most efficient partitioning being 70 percent of the data used for training and the remaining 30 percent used to test the model's performance. As this is a classification problem, we employed five matrices that represented the output as a probability to evaluate how well the three algorithms performed.

The accuracy measures we used to assess how well the algorithm performed were the root mean square error (RMSE), the mean square error (MSE), the mean absolute error (MAE), and the mean absolute error/mean absolute error (MAE/MAE) (MAPE). This is the most used accuracy matrix, which indicates how closely the observed data

align with the projected values of a model, as measured by the model's accuracy. This is referred to as the RMSE. Equation 7 shows the root mean square error (RMSE). Because there are a large number of samples in the dataset, the RMSE is more accurate as a result. We, on the other hand, employed R-squared, which is known as the “coefficient of determination,” as indicated in Equation 8. This is the format that we used.

$$RMS = \sqrt{\frac{1}{n} \sum_{i=1}^n (d_i - p_i)^2} \quad (7)$$

$$R^2 = 1 - \frac{UnexplainedVariation}{TotalVariation} \quad (8)$$

4. Results and discussions

The experimental analysis findings are presented in this section. The dataset comprises 9,999 records, 300 individuals, and 437 locations. A cold start is not an issue in this dataset because all users have given ratings, and all places have been ranked. Figure 3 shows how many people gave each rating. Another thing to note: 1 is the lowest rating category, and 5 is the largest one. For each category, there are 1,706 records, 2,071 records for category 2, 2,096 records for category 3, and 2,106 records for category 4. There are also 2021 records for category 5.

The findings of the evaluation matrices are also provided in Table 2. It is demonstrated that LR has lower RMSE, MSE, MAE, and MAPE values than the other algorithms, resulting in a better fit than the other methods. Aside from that, LR has the highest R-squared value, which means it delivers a superior match as well. It has been noted that NB performs better than DT, but that LR is the optimal algorithm for usage with the presented dataset in general.

Table 2: Summary of Results.

Matrix\ Algorithm	Naive Bayes	Decision Tree	Linear Regression
RMSE	1.823	2.082	1.334
MSE	3.323	4.335	1.779
MAE	1.389	1.544	1.058
MAPE	0.562	0.591	0.467
R-squared	-0.801	-1.349	0.036

As previously described in the preceding section, there are 12 features in total, excluding (User Id, Place Id, and

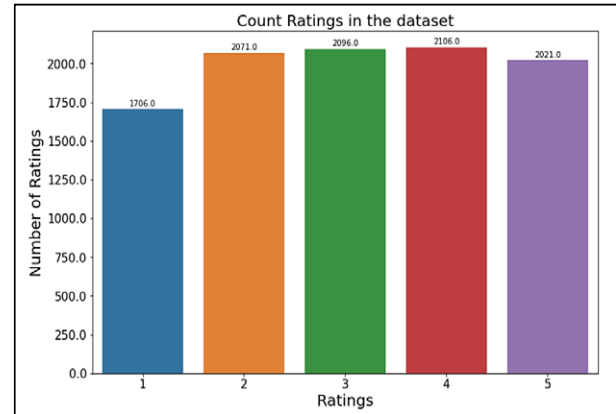


Figure 2. Count the rating in the dataset.

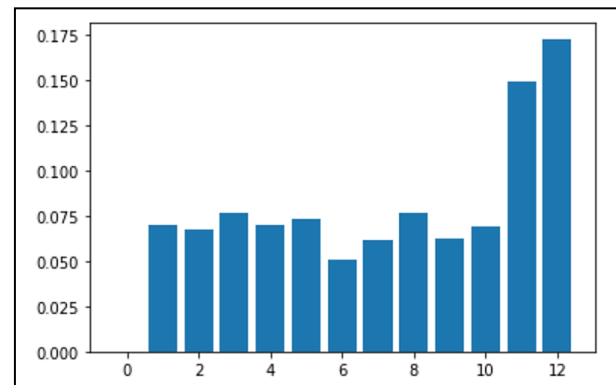


Figure 3. DT feature importance.

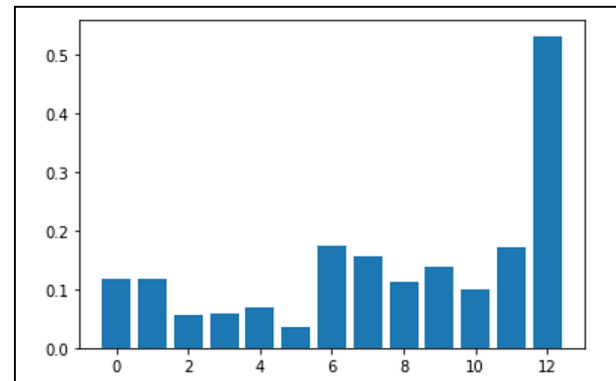


Figure 4. LR feature importance.

rating). Figure 4 illustrates that each characteristic in DT has been assigned a score reflecting its relevance. When examining the relevance of the features, it is noted that feature 0 (global average) has the lowest importance and has been assigned a value of 0.00000. In contrast, feature 12 (place average) has the highest importance and is assigned a score of 0.17275.

Furthermore, as illustrated in Figure 5, each characteristic in LR has been allocated a score that reflects the relevance of that feature. According to the results, it is discovered that feature 5 (similar user rating 5) has the lowest relevance, as indicated by a score of 0.03667, while feature 12 (place average) has the highest value, as indicated by a score of 0.17275.

5. Conclusions

The purpose of this work was to develop a machine learning prediction that can be utilized in an intelligent tourism recommendation system. As was said earlier, recommender systems are a subset of machine learning that serves the purpose of providing users with “suitable” recommendations based on their own interests as well as the interests of others. This study focused on the intelligent tourism recommendation system, and more specifically, on the challenge of providing ideas for tourist attractions, which was ultimately solved. The purpose of this post was to develop an intelligent system that utilizes various machine learning methods to provide suggestions for locations to visit while traveling. A total of three different machine learning algorithms were utilized in the research project: the NB, the DT, and the LR. During the experiment described earlier, we found that the LR model had the best R-squared value, as well as the lowest RMSE, MSE, MAE, and MAPE values. Furthermore, in terms of performance, the NB model outperforms the DT type. On the other hand, the LR model is often considered to be the most suitable model for use with our dataset. In the future, the performance of our suggested models will be enhanced by incorporating a new dataset that includes additional features, as well as other machine learning techniques that we plan to implement in the near future.

Conflict of interest

The authors declare that they have no conflicts of interest.

Funding

The authors received no specific funding for this work.

Acknowledgements

This research is supported by the Department of Computer Science and Artificial Intelligence, College of Computer Science and Engineering, University of Jeddah.

References

- Abbasi-Moud, Z., Vahdat-Nejad, H., & Sadri, J. (2020). Tourism recommendation system based on semantic clustering and sentiment analysis. *Expert Systems With Applications*, 167, 114324.
<https://doi.org/10.1016/j.eswa.2020.114324>
- Al-Ghobari, M., Muneer, A., & Fati, S. M. (2021). Location-Aware Personalized Traveler Recommender System (LAPTA) using Collaborative Filtering KNN. *Computers, Materials & Continua/Computers, Materials & Continua (Print)*, 69(2), 1553–1570.
<https://doi.org/10.32604/cmc.2021.016348>
- Charbuty, B., & Abdulazeez, A. (2021). Classification based on decision tree algorithm for machine learning. *Journal of Applied Science and Technology Trends*, 2(01), 20–28.
<https://doi.org/10.38094/jastt20165>
- Cui, Y., Huang, C., & Wang, Y. (2019). Research on Personalized Tourist Attraction Recommendation based on Tag and Collaborative Filtering. *Proceedings of the 31st Chinese Control and Decision Conference*.
<https://doi.org/10.1109/ccdc.2019.8832961>
- Jeong, C., Lee, J., & Jung, K. (2020). Adaptive Recommendation System for tourism by personality type using deep learning. *International Journal of Internet, Broadcasting and Communication*, 12(1), 55–60.
<https://doi.org/10.7236/ijibc.2020.12.1.55>
- Kim, J., & Lennon, S. (2012). Music and amount of information: do they matter in an online apparel setting? *The International Review of Retail Distribution and Consumer Research*, 22(1), 55–82.
<https://doi.org/10.1080/09593969.2011.634073>
- Kolluri, J., & Razia, S. (2020). WITHDRAWN: Text classification using Naïve Bayes classifier. *Materials Today Proceedings*.
<https://doi.org/10.1016/j.matpr.2020.10.058>
- Mana, S. C., & Sasipraba, T. (2021). A Machine Learning Based Implementation of Product and Service Recommendation Models. *Proceedings of the 7th International Conference on Electrical Energy Systems (ICEES 2021)*, 543–547.
<https://doi.org/10.1109/icees51510.2021.9383732>
- Nilashi, M., Bagherifard, K., Rahmani, M., & Rafe, V. (2017). A recommender system for tourism industry using cluster ensemble and prediction machine learning techniques. *Computers & Industrial Engineering*, 109, 357–368.
<https://doi.org/10.1016/j.cie.2017.05.016>

Olsen, A. A., McLaughlin, J. E., & Harpe, S. E. (2020). Using multiple linear regression in pharmacy education scholarship. *Currents in Pharmacy Teaching and Learning*, 12(10), 1258–1268.

<https://doi.org/10.1016/j.cptl.2020.05.017>

Parveez Saniya, I. R. (2020, October 13). Recommendation system tutorial with Python using collaborative filtering. Towards AI.

<https://pub.towardsai.net/recommendation-system-in-depth-tutorial-with-python-for-netflix-using-collaborative-filtering-533ff8a0e444> (Accessed November 20, 2021)

Patel, K., & Patel, H. B. (2020). A state-of-the-art survey on recommendation system and prospective extensions. *Computers and Electronics in Agriculture*, 178, 105779.

<https://doi.org/10.1016/j.compag.2020.105779>

Ramzan, B., Bajwa, I. S., Jamil, N., Amin, R. U., Ramzan, S., Mirza, F., & Sarwar, N. (2019). An intelligent data analysis for recommendation systems using machine learning. *Scientific Programming*, 2019, 1–20.

<https://doi.org/10.1155/2019/5941096>

Seroussi, Y. (2015). The wonderful world of recommender systems. Yanir Seroussi.

<https://yanirseroussi.com/2015/10/02/the-wonderful-world-of-recommender-systems/> (Accessed October 30, 2021)

Shi, Q., Dong, B., He, T., Sun, Z., Zhu, J., Zhang, Z., & Lee, C. (2020). Progress in wearable electronics/photonics—Moving toward the era of artificial intelligence and internet of things. *InfoMat*, 2(6), 1131–1162.

<https://doi.org/10.1002/inf2.12122>

Umanets, A., Ferreira, A., & Leite, N. (2014). GuideMe – A Tourist Guide with a Recommender System and Social Interaction. *Procedia Technology*, 17, 407–414.

<https://doi.org/10.1016/j.protcy.2014.10.248>

Zhou, X., Su, M., Feng, G., & Zhou, X. (2020). Intelligent Tourism Recommendation Algorithm based on Text Mining and MP Nerve Cell Model of Multivariate Transportation Modes. *IEEE Access*, 9, 8121–8157.

<https://doi.org/10.1109/access.2020.3047264>