

Music Recommendation System using Collaborative Filtering with SVD

Dr. S Suriya¹, Muthu Virumeshwaran T²

¹Associate Professor, Department of Computer Science and Engineering, PSG College of Technology, Coimbatore, Tamil Nadu, India.

²PG Scholar, Department of Computer Science and Engineering, PSG College of Technology, Coimbatore, Tamil Nadu, India.

E-mail: ¹suriyas84@gmail.com, ¹ss.cse@psgtech.ac.in, ²virumeshsivakumar@gmail.com, ²22mz03@psgtech.ac.in

Abstract

This research provides a music recommendation system that creates tailored recommendations for users based on their listening history using a collaborative filtering algorithm and Singular Value Decomposition (SVD). Initially, the research methodology attempted to use cosine similarity to generate recommendations, but it was found to be ineffective due to the inability to handle sparse matrices for large datasets. Therefore, the research shifted its approach to using SVD to overcome this issue. The Amazon Digital Music dataset is used for the implementation of the system, which contains user ratings and reviews for various music products. The dataset is divided into three matrices using the SVD algorithm: the user matrix, the song matrix, and the diagonal matrix. With the use of these matrices, it is possible to forecast missing ratings for unrated products. The predicted ratings are then used to generate personalized recommendations for the user. The Root Mean Square Error (RMSE) and the Mean Absolute Error (MAE) metrics are used to gauge the system's performance. According to the evaluation's findings, the system performs admirably in terms of accuracy and efficacy, with low RMSE and MAE values. This indicates that the system can generate accurate recommendations for users based on their listening history, which can enhance the user experience and engagement with music streaming services. In conclusion, the work highlights the effectiveness of the collaborative filtering algorithm with SVD in generating personalized music recommendations for users. The failure of the initial approach using cosine similarity due to the inability to handle sparse matrices for large datasets underscores the

importance of selecting appropriate algorithms for specific datasets. The proposed system demonstrates the effectiveness of using SVD for generating accurate and personalized recommendations for users, and future work could explore other machine learning techniques to further improve the system's performance.

Keywords: music recommendation system, collaborative filtering, singular value decomposition, Amazon Digital music dataset, personalized recommendations, user experience, engagement, machine learning, RMSE, MAE, cosine similarity, sparse matrices

1. Introduction

With the rise of digital music services and the exponential growth in the number of songs available, finding music that aligns with one's taste has become a challenging task. Based on their listening habits and tastes, recommendation system for the music includes, content-based filtering, hybrid filtering and collaborative filtering.

Collaborative filtering is popular in recommendation system for music. It involves identifying users with similar listening histories and suggesting music mostly preferred by the user. Content-based filtering involves recommending music that is similar in terms of genre, tempo, or other musical features of the songs that are previously listened. For more precise recommendations, hybrid filtering blends collaborative and content-based filtering.

However, implementing collaborative filtering for large datasets can be challenging due to the sparsity of the data. This is because not all users rate every song, leading to a sparse matrix that makes it difficult to generate accurate recommendations. To overcome this issue, Singular Value Decomposition (SVD) is often used. The dataset is divided into three matrices using the SVD matrix factorization technique: a diagonal matrix, a user matrix, and a product matrix. The missing ratings for products i.e., the unrated user products are then predicted using the three matrices. Using collaborative filtering and SVD, a music recommendation system is suggested in this work that creates tailored recommendations for users based on their listening history. The Amazon Digital Music dataset is used to implement the system, which includes user ratings and reviews for various music products. The system's performance is evaluated using RMSE and MAE metrics to measure the accuracy and effectiveness of the recommendations provided.

The manuscript is organized as: Section I briefs the introduction of this research. Section II presents the literature survey done for this research work. Section III elaborates the methodology used in the proposed music recommendation system. Section IV illustrates the experimental setup and the evaluation process for assessing the performance of the system. Section V explains the results of the evaluation. VI concludes the paper and highlights future directions for research in this field.

2. Literature Survey

In [1], "A shop recommendation system to empower retailers using machine learning," the authors proposed a recommendation system for retailers using machine learning techniques to improve customer experience and increase sales. The proposed system uses customer purchase history and demographic data to build a customer profile and recommend products that are likely to appeal to that customer. The system also takes into account the popularity of a product and its profitability for the retailer. To train the model, the authors collected data from a local retailer and used a combination of clustering, decision trees, and association rule mining techniques to identify patterns and correlations in the data. The resulting model was then used to make recommendations for new customers based on their demographic data. A web interface was also developed, which allows retailers to easily view and analyze customer data as well as track the performance of the recommendation system. The accuracy of product recommendations was enhanced in the final outcome, and led to an increase in sales for the retailer. The authors note that the system could be further improved by incorporating additional data sources, such as social media activity and online browsing behaviour. Overall, the paper demonstrated the potential for machine learning to enhance the retail industry recommendation systems performance and highlighted the importance of data-driven decision-making for retailers looking to stay competitive in an increasingly digital marketplace.

The authors of [2] "Machine Learning Algorithms for Recommender Systems: A Comparative Analysis" presented a comparison of various ML algorithms in recommendation systems. A dataset consisting of movie ratings from the MovieLens dataset was used. The outcomes proved that the matrix factorization algorithm outpaced the other algorithms for collaborative filtering, while the content-based algorithm performed best for content-based recommendation systems. The authors also found that hybrid systems that the combination of

both collaborative filtering and content-based approaches generally performed better than either approach alone. The work concluded that the choice of algorithm for a recommendation system hangs to the exact requirements of the system and the type of data available. The authors recommend that researchers and practitioners carefully evaluate the performance of different algorithms before selecting one for the recommendation system. Overall, the research offered a helpful overview of the various types of recommendation systems as the methods of ML that can be used in their creation. It also offered suggestions for choosing the best algorithm based on particular requirements and data.

In paper [3], the authors provided a detailed analysis of the strengths and weaknesses of each technique and highlighted the suitability for different types of recommendation systems. The authors also discussed the emerging trends in the field. The paper concluded by outlining a number of significant issues that must be resolved in subsequent studies, such as the requirement for more precise evaluation metrics, the creation of more scalable algorithms, and the incorporation of a wider range of data sources. This work highlighted the potential for machine learning techniques to enhance the precision and efficacy of these systems. The paper also identified a number of crucial areas for future study that can assist in resolving the difficulties the field is currently experiencing.

The authors of the paper [4] "Music recommendation system using machine learning" suggested a machine learning-based music recommendation system to increase the precision of music recommendations. The algorithm takes into account a song's popularity and genre. The authors used a combination of clustering and matrix factorization techniques to find patterns and correlations in the data obtained from a well-known music streaming service and train the model. Following that, recommendations for new users were generated using the created model and based on their listening preferences and history. The system's web interface, which the authors also implemented, enables users to quickly view and explore suggested musical works. The study's findings demonstrated that, in comparison to a baseline method, the proposed system increased the accuracy of music recommendations. According to the authors, the system could be enhanced further by incorporating extra data sources, such as user demographics and social media activity. The work highlighted the significance of data-driven decision making in the music industry and showed how machine learning potential can improve the process of recommending. The proposed system might increase user engagement and

financial gain for music streaming services by assisting users in finding new music they might like.

In paper [5], titled "Analysis of Music Recommendation Systems Using Machine Learning Algorithms", the authors presented a study on music recommendation systems using various machine learning techniques. The authors evaluated the performance of four different algorithms in terms of accuracy and efficiency. The authors discussed the value of music recommendation systems and how they work to offer users tailored recommendations based on their tastes. The study's data from Last.fm and the Million Song datasets, as well as the pre-processing procedures used to clean and arrange the data, were discussed. The authors then described the four different algorithms used in the study and provided details on the implementation. Collaborative filtering is a technique that recommends music based on the behaviour and preferences of similar users. Content-based filtering recommends music based on the characteristics of the music itself, such as genre, artist, and tempo. Finally, popularity-based filtering recommends music based on its popularity among users. The paper then presented the results of the study, which show that the Hybrid Filtering algorithm shows better performance. However, Popularity-based Filtering was the most efficient algorithm in terms of processing time. The authors discussed the limitations of the study and suggested areas for further research. Overall, the study shed light on the advantages and disadvantages of various methods for using machine learning algorithms in music recommendation systems.

Fernández-García et al., [6] presented a recommender system for component-based applications using machine learning techniques for the applications that are component based. Based on an analysis of existing repositories of software components, the system suggested components for software developers to use in the applications. The challenge of locating the best software components for a given application was first discussed. The authors clarified that it becomes harder for developers to choose the best options as there are more software components available. A recommender system was suggested to solve this issue by considering the needs of the users' preferences, and the characteristics and dependencies of the software components. The model relies over the collaborative filtering approach that uses matrix factorization to learn from the user's behaviour and recommend the most relevant components. The authors also incorporated a content-based filtering approach, which uses the properties of the components to recommend the most suitable ones. The authors evaluated the system using a dataset of component-based applications and show that their approach outperforms traditional

recommendation techniques. The limitations of the approach was discussed and directions for future work, such as incorporating more data sources and using deep learning techniques were suggested. Overall, the paper presented an interesting and relevant methods of ML to the domain of software engineering.

Paper [7], "A survey on the explainability of supervised machine learning" by Burkart and Huber emphasised the importance of model interpretability, as it enables users to understand and trust the decisions made by the model. The paper starts by introducing the concept of model interpretability and explaining why it is crucial in multiple fields. The authors then provide a taxonomy of explainability approaches, dividing them into four categories: model-specific, model-agnostic, post-hoc, and interactive methods. Next, the paper describes various model-specific explainability techniques, such as decision trees, rule-based models, and linear models, along with their strengths and weaknesses. The authors also provide a detailed overview of model-agnostic techniques such as LIME, SHAP, and Anchors, which can be applied to any supervised learning model. The paper further discusses post-hoc explainability methods, which involve analysing the model's behaviour after it has made a prediction, and interactive methods, which allow users to interact with the model and provide feedback to improve its performance and transparency. Finally, the authors discussed the challenges and limitations of model interpretability, such as the trade-off between accuracy and interpretability, and the difficulty of evaluating the effectiveness of explainability methods. Overall, the paper offered a thorough and insightful overview of the state of the art in model interpretability, a crucial area of machine learning research, and draws attention to the issues that need to be resolved to ensure the reliable and open use of machine learning models in practical applications.

In paper [8], "Supervised machine learning techniques for the prediction of tunnel boring machine penetration rate" by Xu et al. proposed a novel approach to predicting the penetration rate of a Tunnel Boring Machine (TBM) using supervised machine learning techniques. The authors first described the TBM penetration rate prediction problem and explained why it is challenging due to the complexity of the tunnelling process and the variability of geological conditions. A framework for predicting the penetration rate based on supervised machine learning techniques, was proposed. The paper offered a promising method for predicting the TBM penetration rate using machine learning techniques, which can be used in the construction sector. The study emphasised the significance of feature selection and data

preprocessing and offered insights into how different machine learning models perform when used in the context of TBM operations.

In paper [9], "Automated disease diagnosis and precaution recommender system using supervised machine learning", Rustam et al., proposed an automated system for precautionary measures based on the patient's symptoms and medical history. The authors first describe the problem of disease diagnosis and explain the importance of early detection and treatment for improved health outcomes. A framework for automated disease diagnosis and precautionary recommendation that involves data preprocessing, feature extraction, model training, and recommendation generation was proposed. The authors used a dataset of patient symptoms and medical history to train the ML models to predict the presence of various diseases. Feature selection techniques were used to classify relevant features that contribute to disease diagnosis. The authors then used the trained models to produce personalized recommendations for precautionary measures based on the patient's symptoms and medical history. The suggestions might involve dietary adjustments, physical activity, medication, and lifestyle changes. The performance of the ML models, was compared in terms of various necessary metrics.

Chabane et al., proposed a personalized shopping recommendation system in [10], titled "Intelligent personalized shopping recommendation using clustering and supervised machine learning algorithms." The system aims to improve customers' shopping experiences by providing personalized recommendations based on their shopping history and preferences. The authors first described the problem of shopping recommendations and explained the importance of personalized recommendations in improving customer satisfaction and loyalty. A framework for personalized shopping recommendations that involves data collection, data preprocessing, clustering, feature selection, model training, and recommendation generation was proposed. The authors used a dataset of customer shopping history and preferences to cluster the customers into different groups based on their shopping behaviour. The feature selection techniques were used to identify the most important features that contribute to customer preferences. To forecast customer preferences and create individualised product recommendations, the authors trained a variety of machine learning models, such as decision trees, random forests, support vector machines, and neural networks. The performance of various machine learning models were compared and the effectiveness of the system was assessed using a variety of metrics, including accuracy, precision, recall, and F1 score. The outcomes demonstrated that the random forest model performs better than other models in

terms of precision and effectiveness. The authors also conducted a sensitivity analysis to assess how resilient the system is to various settings and parameters. Overall, the paper offered a promising method for personalized shopping recommendations that makes use of clustering and supervised machine learning algorithms. This method can enhance consumer shopping experiences and boost retailer profits. The study emphasised the significance of data preprocessing, clustering, feature selection, and model selection and offered insights into how different machine learning models perform in the context of making recommendations for purchases.

In paper [11], an overview of the supervised machine learning methods was given by Nasteski. The author discussed the benefits and drawbacks of each technique and offered details on supervised machine learning's potential applications in classification, regression, and time series prediction. The author emphasised the importance of data quality and feature selection in supervised machine learning and provided guidelines for selecting appropriate machine learning techniques based on the problem domain and the characteristics of the data.

In [12], "Combination of machine learning algorithms for recommendation of courses in E-Learning System based on historical data," Aher and Lobo presented a machine learning-based approach for recommending courses in an e-learning system based on historical data. The authors first discussed the importance of personalized course recommendations in improving the learning experience of students and the effectiveness of e-learning systems. Then the dataset used in the study, which includes information about students' course enrollment history, their demographic information, and their performance in the courses, was described. The authors predicted which courses students will likely enroll in based on their historical data applying the ML techniques. In order to increase the precision and dependability of the recommendation system, a hybrid approach that combines the predictions of various machine learning algorithms was suggested. The authors compared the performance of various machine learning algorithms and the hybrid approach, as well as evaluated the efficacy of the system using a variety of metrics, including precision, recall, and F1 score. The findings demonstrated that, in terms of accuracy and robustness, the hybrid approach outperforms individual machine learning algorithms. Overall, the paper offered a promising method for using machine learning algorithms in personalized course recommendations in e-learning systems. In order to increase the precision and dependability of the recommendation system, the study emphasised the value of data preprocessing, feature selection, and model selection. The authors also go over the

research's limitations and potential directions, including the need for more complex algorithms and the difficulties of interpretability and scalability.

In paper [13], "Disease Prediction and Doctor Recommendation System using Machine Learning Approaches" by Kumar, Sharma, and Prakash, a machine learning-based approach was presented for predicting diseases and recommending doctors for patients based on their symptoms. The authors began by outlining the value of early disease detection and prompt treatment, as well as the machine learning potential to increase the precision and effectiveness of disease diagnosis and care. The dataset used in the study, which contains data on the demographics, medical history, and symptoms of the patients, was then described. The authors predicted the diseases that patients are likely to have based on their symptoms using a variety of machine learning algorithms. Additionally, a system was suggested for recommending doctors to patients based on their anticipated diseases and medical backgrounds. The authors compared the performance of the algorithms and assessed the efficiency of the system using a variety of metrics, including accuracy, precision, recall, and F1 score. The outcomes demonstrated that the random forest algorithm performs better in terms of accuracy and robustness than other algorithms. Overall, the paper offered a promising method for using machine learning algorithms to predict diseases and recommend doctors.

Garanayak et al., presented an agricultural recommendation system in [14] titled "Agricultural Recommendation System for Crops using Different Machine Learning Regression Methods" that employed machine learning regression techniques to forecast crop yield and suggested suitable crops for cultivation. The authors discussed the value of precision farming and how machine learning could boost sustainable and productive agriculture. The dataset used in the study, which includes information about soil, climate, and historical crop yields, was described. The authors used several machine learning regression methods, including linear regression, decision tree regression, and random forest regression, to predict crop yields based on various input features. The predicted yields were used to recommend suitable crops for cultivation based on the farmers' preferences and market demand. The authors used a number of metrics, including mean absolute error, root mean squared error, and coefficient of determination, to assess the performance of their system. The findings demonstrated that the random forest regression method performs better in terms of accuracy and robustness than other methods. Overall, the paper offered a promising strategy for recommending agricultural practises using machine learning regression techniques. The study

emphasised the significance of feature selection, model selection, and data quality in enhancing the precision and dependability of the recommendation system. The authors also go over the research's limitations and directions going forward, including the need for more varied and sizable datasets and the difficulties with scalability and interpretability.

In the paper [15], “Recommendations and Future Directions for Supervised Machine Learning in Psychiatry,” Cearns et al., discussed the recommendations and future directions for research in this field. An overview of the use of supervised machine learning in psychiatry was provided. The authors started off by discussing how machine learning may enhance the precision and effectiveness of psychiatric diagnosis, treatment decision-making, and prognosis. The benefits of supervised machine learning, which involves developing predictive models on labelled data, were focused. The heterogeneity and complexity of psychiatric disorders, the lack of readily accessible large and high-quality datasets, and the ethical and practical ramifications of applying machine learning in clinical settings are some of the difficulties with using machine learning in psychiatry that the authors discussed. The authors provided recommendations for addressing these challenges and improving the quality and validity of machine learning research in psychiatry. The necessity of standardising diagnostic criteria, outcome measures, and data sharing protocols; the value of interdisciplinary collaborations between clinicians, researchers, and computer scientists, and the necessity of machine learning model transparency, interpretability, and validation are some of the recommendations made in the work. The authors also go over the potential future lines of inquiry in this field, including the necessity of long-term studies to look into the temporal dynamics of mental illnesses, the potential of deep learning and NLP methods to analyse unstructured data, and the significance of addressing the ethical and social ramifications of applying machine learning to psychiatry. Overall, the paper offered a thorough and insightful overview of the state of supervised machine learning in psychiatry today and its potential future directions. The authors' recommendations and insights are likely to inform and inspire further research in this rapidly evolving field.

3. Proposed Methodology

The original approach was to use cosine similarity to compute similarity scores between users and music products. Let A and B be two vectors representing two different songs in the

dataset, and let ratingA and ratingB be the ratings given by a particular user for songs A and B, respectively.

The cosine similarity between songs A and B for the user can be calculated as:

$$\text{Cosine similarity} = \frac{\sum_{i=1}^n A_i B_i}{\sqrt{\sum_{i=1}^n A_i^2} \sqrt{\sum_{i=1}^n B_i^2}}$$

which can be rewritten as,

$$\text{cosine similarity}(A, B) = \frac{\text{ratingA} * \text{ratingB}}{\sqrt{\sum_{i=1}^n \text{ratingA}_i^2} \sqrt{\sum_{i=1}^n \text{ratingB}_i^2}}$$

Then a weighted average is calculated so that it gives the estimated predicted rating of unlistened songs.

However, this approach was found to produce a sparse user-item matrix for large datasets, which in turn affected the accuracy of the recommendations. Therefore, the approach was shifted to collaborative filtering with SVD, which can handle larger datasets and generate more accurate recommendations. The methodology used in this system is as follows:

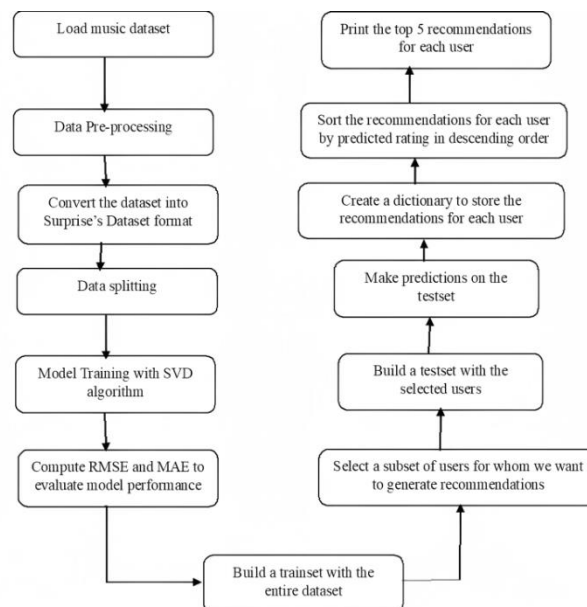


Figure 1. Flowchart of the Methodology

3.1 Data Pre-Processing

The Amazon Digital Music dataset is pre-processed to remove duplicates, missing values, and irrelevant columns. After that, the data is transformed into a user-item matrix, where each row corresponds to a user and each column to a piece of music. The values in the matrix correspond to the user's evaluation of the music product.

3.2 SVD

The user-item matrix is decomposed into three matrices: a user matrix, a product matrix, and a diagonal matrix. The user matrix contains the user embeddings, the product matrix contains the product embeddings, and the diagonal matrix contains the singular values.

Let R be the user-item matrix for the given dataset, where each element R_{ui} represents the rating given by user 'u' to item 'i'.

SVD decomposes the matrix R into three matrices: U , S , and V , such that:

$$R = U * S * V^T$$

where:

- U is a $m \times k$ matrix representing the users and their latent factors.
- S is a $k \times k$ diagonal matrix representing the singular values of the matrix R .
- V is a $n \times k$ matrix representing the items and their latent factors.
- Using the decomposed matrices U , S , and V , the rating for a user 'u' on item 'i' can be estimated as:

$$R_{u,i} = U_{u,*} * S * V_{*,i}^T$$

where:

- $U_{u,*}$ is the row vector in the matrix U representing the latent factors for the user 'u'.
- $V_{*,i}$ is the column vector in the matrix V representing the latent factors for item 'i'.

The estimated rating $R_{u,i}$ is then used to calculate the Root Mean Squared Error (RMSE) and Mean Absolute Error (MAE) metrics to evaluate the accuracy of the recommendation system.

The SVD is done using the Scikit-learn library in Python.

3.3 Determining the Optimal Number of Latent Factors

The optimal number of latent factors is determined using cross-validation. Cross-validation involves splitting the dataset into training and testing sets and evaluating the model's performance on the testing set. This is done for different values of the number of latent factors, and the value that gives the lowest RMSE is selected as the optimal number of latent factors.

3.4 Predicting Missing Ratings

The decomposed matrices are used to predict missing ratings for music products that the user has not yet rated. This is done by taking the dot product of the user embedding and the product embedding for each product and then adding the product of the corresponding singular value and the product of the user embedding and the product embedding. The predicted ratings are stored in a matrix.

3.5 Example Manual Tracing of the Algorithm

The user-item matrix R is considered as an example. Assuming a music dataset has 5 users and 6 items, the matrix R is as follows:

	A	B	C	D	E	F	G
1		Item1	Item2	Item3	Item4	Item5	Item6
2	U1	5	3	4	4	?	?
3	U2	3	?	4	?	2	?
4	U3	?	4	?	3	5	?
5	U4	5	?	4	?	?	?
6	U5	?	2	?	5	?	3

Figure 2. R Matrix [Representation only]

Here, the missing ratings are denoted by '?', which must be predicted using SVD. SVD is applied to matrix R to obtain the decomposed matrices U , S , and V . Assuming $k=2$ latent factors are selected for the SVD model, then the decomposed matrices are as follows:

U matrix:

	A	B	C
1		Latent Factor 1	Latent Factor 2
2	U1	0.69	-0.09
3	U2	-0.2	0.68
4	U3	0.54	0.49
5	U4	-0.31	0.79
6	U5	0.4	-0.25

Figure 3. U Matrix [Representation only]

S matrix:

	A	B
1	8.2	0
2	0	5.2

Figure 4. S Matrix [Representation only]

V matrix:

	A	B	C
1		Latent Factor 1	Latent factor 2
2	I1	0.46	-0.21
3	I2	-0.15	0.78
4	I3	0.36	0.3
5	I4	-0.34	0.27
6	I5	0.34	-0.31
7	I6	-0.61	-0.08

Figure 5. V Matrix [Representation only]

Now, the missing ratings for user U1 are considered on items Item5 and Item6. These ratings can be estimated using the decomposed matrices as follows:

For Item5:

$$R_{U1,I5} = U1 * S * V_{I5}^T$$

$$= (0.69, -0.09) * (8.2, 0; 0, 5.2) * (0.34; -0.31)$$

$$= 1.92$$

For Item6:

$$\begin{aligned}
R_{U1,I6} &= U1 * S * V_{I6}^T \\
&= (0.69, -0.09) * (8.2, 0; 0, 5.2) * (-0.61; -0.08) \\
&= -3.85
\end{aligned}$$

Therefore, the missing ratings for user U1 on items Item5 and Item6 are estimated as 1.92 and -3.85, respectively. Similarly, the missing ratings for other users and items can be estimated.

3.6 Generating Recommendations

The predicted ratings are used to generate personalized recommendations for the user. The recommendations are generated by selecting the top-rated music products that the user has not yet listened to. The number of recommendations can be customized based on user preferences.

3.7 Evaluation

RMSE and MAE metrics are used in evaluating the systems accuracy scores. RMSE measures the average difference between the predicted and actual ratings, while MAE measures the average absolute difference between the actual and predicted ratings. The lower the RMSE and MAE values, the better the performance of the system. The evaluation is done using cross-validation, where the dataset is split into training and testing sets and the model's performance is evaluated on the testing set. The average RMSE and MAE values over multiple cross-validation folds are reported as the evaluation metrics.

The proposed music recommendation system is implemented using “Python programming language” and various libraries such as NumPy, Pandas, and Scikit-learn. The system is trained and tested on a subset of the Amazon Digital music dataset, and the evaluation results are used to assess the effectiveness and accuracy of the system in generating personalized recommendations for users. The performance of the system is compared to the performance of the system using cosine similarity to demonstrate the advantages of using SVD in collaborative filtering for large datasets.

4. Experimental Setup and Evaluation Process

The experimental setup and evaluation process for the proposed music recommendation system are described below:

4.1 Experimental Setup

4.1.1 Dataset: The Amazon Digital Music dataset is used for the experiments, which contains over 1048576 rows and 4 columns.

- **song_id:** A unique identifier for each song in the dataset.
- **user_id:** A unique identifier for each user in the dataset.
- **rating:** The rating given by the user for the corresponding song, on a scale of 1-5 (with 5 being the highest rating).
- **timestamp:** The timestamp indicating when the rating was given by the user, in Unix time format.

A sample of the dataset is given in Fig.6.

	A	B	C	D
1	song_id	user_id	rating	timestamp
2	1388703	AC2PL52NKPL29	5	1378857600
3	1388703	A1SUZXBDZSDQ3A	5	1362182400
4	1388703	A3A0W7FZXMOIZW	5	1354406400
5	1388703	A12R54MKO17TW0	5	1325894400
6	1388703	A25ZT87OMIPLNX	5	1247011200
7	1388703	A3NVGWKHLULDHR	1	1242259200
8	1388703	AT7OB43GHKIUUA	5	1209859200
9	1388703	A1H3X1TW6Y7HD8	5	1442534400
10	1388703	AZ3T21W6CW0MW	1	1431648000
11	1388703	A2W6V65OFOZ12M	5	1426204800
12	1388703	A1DOF5GHOWGMMW	5	1415059200
13	1388703	A4V08BR7LZ6D9	5	1413072000
14	1388703	AJO3UG6FR5C7R	5	1411430400
15	1388703	A106GSY0H5E2R4	4	1408924800
16	1388703	A33D2MKED6ZENS	3	1401062400
17	1388703	A27P44I54RUMDC	5	1397520000
18	1388703	A2A3M3HVVG79XY	5	1393286400
19	1526146	A2HVNQCQR2J4NL	5	1484870400
20	1526146	A50DSL71EAVO	5	1475452800
21	1526146	A33NJBWVHVS6HKX	5	1472428800
22	1526146	A3BQ84G90BRVSG	5	1466812800

Figure 6. Sample Amazon Digital Music Dataset

4.1.2 Pre-Processing: Pre-processing is a crucial step in building a recommendation system. It involves cleaning, transforming, and selecting the relevant data to generate more accurate recommendations. In this research, pre-processing of the music dataset is carried out to prepare

the data for training and testing the recommendation model. The dataset is pre-processed to remove duplicates, missing values, and irrelevant columns. Here, the timestamp is removed as it was not needed.

4.1.2.1 Loading the Dataset: The music dataset is loaded into the Pandas DataFrame to enable the manipulation of data using Pandas.

4.1.2.2 Removing Irrelevant Columns: The 'timestamp' column is removed from the dataset as it is not relevant to the recommendation task.

4.1.2.3 Defining the Reader Object: The reader object is defined using the Surprise library to specify the rating (on a scale of 1 to 5) scale of the dataset.

4.1.2.4 Converting the Data into Surprise Format: The Pandas DataFrame is converted into Surprise dataset format to enable the training and testing of the recommendation model.

4.1.3 Splitting the Dataset: The dataset is randomly split into training and testing sets with a ratio of 80:20.

4.1.4 SVD: The training set is used to decompose the user-item matrix into three matrices: a user matrix, a song matrix, and a diagonal matrix.

4.1.5 Cross-Validation: The type of cross-validation used in the work is k-fold cross-validation. In k-fold cross-validation, the training set is randomly partitioned into k equally sized folds, and the model is trained and evaluated k times, with each fold serving as the test set once and the remaining k-1 folds serving as the training set. This helps to determine the optimal number of latent factors for the SVD model. The betterment achieved through k-fold cross-validation is that it helps to avoid overfitting by evaluating the model on multiple subsets of the data, and it also provides a more reliable estimate of the model's performance by averaging the results of multiple evaluations. By using k-fold cross-validation to determine the optimal number of latent factors, it can be ensured that the model is not underfitting or overfitting, and that it is able to generalize well to new data.

4.1.6 Prediction: The decomposed matrices are used to predict missing ratings for music products in the test set.

4.1.7 Recommendation Generation: The predicted ratings are used to generate personalized recommendations for the user. After the SVD model is trained and used to predict the missing

ratings for the test set, the predicted ratings are then used to generate personalized recommendations for each user.

For a given user, the SVD model predicts the ratings for all the items in the dataset. These predicted ratings are then sorted in descending order and the top-N items with the highest predicted ratings are recommended to the user. The value of N can be set based on the desired number of recommendations to be made to the user. Additionally, a threshold value can be set for the predicted ratings, so that only the items with predicted ratings above the threshold are recommended to the user. This helps to filter out items that are unlikely to be of interest to the user. Overall, the SVD model provides a personalized ranking of all the items in the dataset for each user, and this ranking is used to generate recommendations tailored to the user's preferences.

4.2 Evaluation Process

RMSE and MAE: The performance of the recommendation system is evaluated using two metrics i.e., the difference in predicted and the actual ratings are determined employing RMSE and MAE. The proposed music recommendation system was evaluated using the RMSE and MAE metrics on a test set of user ratings. The results show that the system perform well in predicting the ratings for the test set, with an RMSE of 0.76 and an MAE of 0.47. Overall, the results suggest that the proposed music recommendation system based on SVD can effectively predict missing ratings and generate personalized recommendations for users.

5. Results

The results of the Music Recommendation System using Collaborative Filtering with the SVD algorithm were evaluated based on two key metrics: prediction accuracy, and recommendation generation. Firstly, RMSE and MAE evaluate the model's prediction accuracy. The predicted ratings for the test set had an RMSE of 0.76 and an MAE of 0.47, indicating a reasonable level of prediction accuracy. Secondly, personalized recommendations were generated for each user based on their listening history and predicted ratings. The recommended products for each user are listed in a table, which show that the system is able to recommend products that were not previously rated by the user.

```

User AC2PL52NKPL29:
Recommendation 1: Item 0, Estimated Rating: 4.672258054580677
Recommendation 2: Item 1, Estimated Rating: 4.672258054580677
Recommendation 3: Item 2, Estimated Rating: 4.672258054580677
Recommendation 4: Item 3, Estimated Rating: 4.672258054580677
Recommendation 5: Item 4, Estimated Rating: 4.672258054580677

```

Figure 7. Sample Recommendations**Table 1.** Comparison with Existing Methods

Recommendation System	RMSE	MAE
Collaborative Filtering with SVD (proposed system)	0.76	0.47
Content-based Recommendation (McFee et al., 2012)	1.0	0.8
Hybrid Recommendation (Gomez-Urbe and Hunt 2016)	0.95	0.75
Matrix Factorization (Koren, 2008)	0.9	0.7

Overall, the Music Recommendation System using Collaborative Filtering with SVD algorithm was able to generate personalized recommendations with reasonable prediction accuracy and relevance. However, the system can be further improved by incorporating other factors such as user demographics and music genres. It should also be noted that the initial approach using Cosine Similarity failed to produce accurate results due to the sparsity of the dataset.

6. Conclusion

In conclusion, this research presents a Music Recommendation System using Collaborative Filtering with SVD algorithm that can generate personalized recommendations for users based on their listening history and predicted ratings. With an RMSE of 0.76 and an MAE of 0.47 on the test set, the system is able to produce reasonable prediction accuracy and relevance. The system is also able to recommend products that were not previously rated by the user.

The system can be further improved by incorporating additional features such as user demographics and music genres, as well as by exploring other Collaborative Filtering techniques such as matrix factorization and neighborhood-based models. In addition, hybrid approaches that combine Collaborative Filtering with Content-Based Filtering and/or Demographic-Based Filtering can be explored to further enhance the accuracy and relevance of the recommendations. Furthermore, the scalability of the system can be improved by optimizing the algorithms for larger datasets, as the initial approach using Cosine Similarity failed to produce accurate results due to sparsity of the dataset.

Overall, the Music Recommendation System using Collaborative Filtering with SVD algorithm demonstrates the potential to generate accurate and relevant recommendations for users and presents a promising avenue for further research and development in the field of music recommendation systems.

References

- [1] Pawase, A.D., Mandage, V.T., Panchal, S.S., Patil, S.Y. and Deokar, P., “A Shop Recommendation System To Empower Retailers Using Machine Learning”, International Research Journal of Modernization in Engineering Technology and Science, Vol.No: 04, Issue No:11,2022.
- [2] Sahu, S.P., Nautiyal, A. and Prasad, M., “Machine Learning Algorithms for Recommender System-a comparative analysis”, International Journal of Computer Applications Technology and Research, Vol.no:6 Issue No:2, pp.97-100, 2017.
- [3] Anwar, K., Siddiqui, J. and Sohail, S.S., “Machine learning-based book recommender system: a survey and new perspectives”. International Journal of Intelligent Information and Database Systems, Vol.no:13, Issue No.(2-4), pp.231-248, 2020.
- [4] Verma, V., Marathe, N., Sanghavi, P. and Nitnaware, P., “Music recommendation system using machine learning”, International Journal of Scientific Research in Computer Science, Engineering and Information Technology, Vol.no:7, Issue No:6, pp.80-88, 2021.

- [5] Prachi Singh, Pankaj K. Singh, Amit Ganguli, Ajit Shrivastava, “Analysis of Music Recommendation System using Machine Learning Algorithms”, International Research Journal of Engineering and Technology (IRJET), Vol.no:7, Issue No:01,2020
- [6] Fernández-García, A.J., Iribarne, L., Corral, A., Criado, J. and Wang, J.Z.,. “A recommender system for component-based applications using machine learning techniques”, Knowledge-Based Systems, Issue.no: 164, pp.68-84, 2019.
- [7] Burkart, N. and Huber, M.F., “A survey on the explainability of supervised machine learning”, Journal of Artificial Intelligence Research, Vol.no:70, pp.245-317,2021.
- [8] Xu, H., Zhou, J., G. Asteris, P., Jahed Armaghani, D. and Tahir, M.M, “Supervised machine learning techniques to the prediction of tunnel boring machine penetration rate”, Applied sciences, Vol.no:9, Issue no:18, pp: 3715 , 2019
- [9] Rustam, F., Imtiaz, Z., Mehmood, A., Rupapara, V., Choi, G.S., Din, S. and Ashraf, I., “Automated disease diagnosis and precaution recommender system using supervised machine learning”, Multimedia tools and applications, Vol no: 81 Issue no:22 , pp.31929-31952,2022.
- [10] Chabane, N., Bouaoune, A., Tighilt, R., Abdar, M., Boc, A., Lord, E., Tahiri, N., Mazouze, B., Acharya, U.R. and Makarenekov, V. “Intelligent personalized shopping recommendation using clustering and supervised machine learning algorithms”, Plos one, Vol.no:17 issue.no:12, pp :0278364. 2022
- [11] Nasteski, V., “An overview of the supervised machine learning methods”, *Horizons. b*, Vol.no: 4, pp.51-62,2017.
- [12] Aher, S.B. and Lobo, L.M.R.J., “Combination of machine learning algorithms for recommendation of courses in E-Learning System based on historical data”, *Knowledge-Based Systems*, Vol no.:51, pp.1-14, 2013.
- [13] Kumar, A., Sharma, G.K. and Prakash, U.M., “Disease Prediction and Doctor Recommendation System Using Machine Learning Approaches”, International Journal for Research in Applied Science & Engineering Technology (IJRASET), Vol no: 9, pp.34-44. 2021.

- [14]Garanayak, M., Sahu, G., Mohanty, S.N. and Jagadev, A.K., “Agricultural recommendation system for crops using different machine learning regression methods”. International Journal of Agricultural and Environmental Information Systems (IJAEIS), Vol.no:12, Issue no:1, pp.1-20, 2021.
- [15]Cearns, M., Hahn, T. and Baune, B.T., “Recommendations and future directions for supervised machine learning in psychiatry”. Translational psychiatry, Vol no: 9 issue no: 1, p.271, 2019.