

Music Recommendation Systems using Exponential Decay Moment based Long Short-Term Memory

Snehal P. Dongre¹ Allabakash Isak Tamboli²

¹Department of Computer Science and Engineering, Yeshwantrao Chavan College of Engineering, Hingna Road, Wanadongri, Nagpur 441110

²Department of Electronics and Telecommunication Engineering, Shri Guru Gobind Singhji Institute of Engineering and Technology, Vishnupuri Nanded-431606, MS, INDIA

*Corresponding Author
Snehal P. Dongre

Article History

Received: 08.07.2025

Revised: 16.08.2025

Accepted: 25.09.2025

Published: 22.10.2025

Abstract: Music recommendation systems explore the fundamental relationship between the users and music to deliver personalized recommendations. However, distinguishing between the preferences of many users with similar music tastes is challenging for convolutional methods. Refining fine-grained music features is a challenging task, often resulting in poor performance in music recommendation systems. To address this challenge, this research aims to develop a novel Deep Learning (DL) approach for effectively recommending music. Therefore, this research proposes the Exponential Decay Moment based Long Short-Term Memory (EDM-LSTM) approach for the recommendation of music. The use of EDM supports LSTM networks by improving the training process and enhancing the effectiveness of the recommendation system through the efficient management of gradient updates during training. Initially, this research uses the two standard music datasets, Lastfm and 30Music to evaluate the effectiveness of the proposed method. This research involves three parts namely music coding, user coding and prediction. The experimental results show that the proposed EDM-LSTM attains the better Recall@10 of 22.370 and 33.303 on Lastfm and 30Music datasets as compared to the existing method like Graph based Attentive Sequential Model with Metadata (GASM).

Keywords: Deep learning, Exponential decay moment, Fine-grained features, Long Short-Term Memory, Music recommendation and User coding.

INTRODUCTION

Recommendation systems have enhanced music streaming platforms by personalizing user experiences and recommending tracks based on individual preferences [1]. These systems improve music finding by delivering relevant songs and artists at the right time, making the listening experience more enjoyable [2][3]. Music is a globally cherished art form, that continues to evolve in both diversity and global reach, playing an important role in the daily lives of various people.

Successful music firms have given rise to various international brands and online music platforms, such as iTunes and PolyGram, among others [4][5]. However, selecting the appropriate songs for the user based on the listening history can be challenging, especially when the user's interests or profile are unclear, and there is inadequate interaction data available [6]. Hence, recommendation systems are the most effective information filtering tool that helps to overcome this problem of information overload [7]. The recommender excels at capturing users' personal preferences, even when dealing with large volumes of data, addressing a critical need in the music recommendation landscape [8]. With the rapid growth of digital content online and the constant rise of streaming platforms, music recommendation has evolved into a highly complex task

and a valuable tool for enhancing user experience [9] [10]. Music recommendations on streaming platforms is inspiring as compared to other recommendation platforms such as books, movies, tourism and more [11]. This is due to the high consumption rates and the vast number of available items, which play a significant role in recommendations [12]. Therefore, scalability is a central requirement for music recommendation approaches. Traditional Machine Learning (ML) and Deep Learning (DL) approaches have been used in the recommendation process, but their shallow structure limits the learning of music features, preventing efficient feature extraction feature [13][14]. Manual extraction of the music features is limited and lacks robustness. Therefore, modern DL methods have been employed to identify deep music features for the recommendation process [15]. However, conventional methods struggle to differentiate between the diverse preferences of users for similar music and to capture the fine-grained features of music. These challenges often resulted in poor performance in recommending music. Hence, this research proposes the Exponential Decay Moment based Long Short-Term Memory (EDM- LSTM) approach for music recommendation. The foremost contributions of this research are as follows: This research aims to propose the EDM- LSTM approach for the recommendation of

music. The EDM allows the recommender to efficiently capture long-term dependencies in the sequential data of user listening behaviors.

LSTM is effective at learning patterns in sequences, allowing for a better understanding of the music users listen to. This improves the performance of future music recommendations by taking historical preferences into account over different periods. This research proposes the model based on three significant parts such as music coding, user coding and prediction. These three components effectively learn the music representations based on the listening behaviors of the users. This research paper is organized as follows: Section 2 presents the literature review. Section 3 presents the proposed methodology and Section 4 gives the results and discussion. The conclusion of this research paper is given in Section 5.

Literature Survey

Here, the existing works related to DL-based music recommendation systems are discussed, along with their advantages and limitations.

Kumar and Kumar [16] presented multiple Machine Learning (ML) approaches for the music recommendation process. The Popularity Predictor (PoP), Sequential pattern mining, Markov decision process and various Nearest Neighbour based ML approaches have been developed and applied to the datasets from various domains. The presented approach effectively identified the frequent patterns and sequences in user interactions and effectively captured the temporal dynamics. However, the multiple ML approach encounters difficulties with sparse user-item interaction matrices, leading to unreliable similarity measures.

Tao [17] developed the Agnostic Inference and Representation Ensemble (AIRE) approach for the recommendation of music. The developed AIRE approach effectively learned the representation of the base recommenders from the interaction history and memorized the characteristics as well as patterns with those representations. The benefit of few shots learning was utilized in this method to learn the base recommender representation with the help of prototype network. Nevertheless, AIRE heavily depends on the input data quality, inaccurate or noisy data leads to unreliable recommendations.

Tofani [18] introduced multiple approaches like Information Retrieval based Term Frequency and Inverse Document Frequency (IR-TFIDF) and IR based Neural Network (IR 1NN) and the third method was IR based Markov Model (IR MC). The developed approach adapted to the specific characteristics of the data without the requirement for hand-engineered features. However, the IR-based approach was not designed to accommodate changes in the relevance of items over time, impacting its ability to handle time-awareness effectively. So, it does not inherently address the temporal dynamics of user preferences.

Weng [19] developed the Graph based Attentive Sequential Model with Metadata (GASM) for the MRS. The developed approach integrated metadata to improve user behavior representation and significantly identified the listening patterns. Initially, a directed listening graph was used to model the relationships between different types of nodes. Then, a Graph Neural Network (GNN) was performed to learn the latent representation vectors. GNNs provided personalized recommendations by learning user embeddings based on their interactions with items. However, GASM was vulnerable to overfitting by allowing it to memorize noise and specific details to the training data, so it leads to poor performance recommendations.

Lu [20] implemented the Multilayer Attention Representation (MAR) for music recommendation. The implemented approach learned the representations of the songs from multidimensions with the help of user-attribute and song preference data, enabling it to extract the relationships among the songs and users effectively. The feature-dependent attention network was developed to capture the variations in user favorites for various historical behaviors. Then, the developed the song-dependent attention network to identify the temporal dependence of the behaviors of the users. However, the MAR model required large volumes of high-quality data to learn meaningful attention patterns.

From this overall analysis, some of the limitations have been identified: difficulty with sparse user-item interaction matrices. Depending only on the quality of input data quality, inaccurate or noisy data leads to unreliable recommendations. The IR-based approach was not prepared to maintain the changes in the relevance of the items over time. Due to the complex architecture of LSTM networks, they are vulnerable to overfitting, particularly when training data is constrained. Thus, large volumes of high-quality data are required to learn meaningful attention patterns. To solve these problems, the EDM-LSTM approach is proposed to recommend the music by capturing the temporal dependencies in a user's music listening history. Through effectively capturing temporal dependencies, EDM-LSTM has effectively understood a user's listening patterns over time, leading to more accurate music recommendations.

Proposed Methodology. This research proposes the EDM-LSTM approach for the recommendations of music systems. The EDM mechanism helps manage the volume of data that the network needs to process, making the approach more scalable for large datasets and user bases. In this research, the two standard music datasets Lastfm and 30Music are used to estimate the effectiveness of the proposed EDM-LSTM approach. Fig. 1 depicts the framework of the EDM-LSTM.

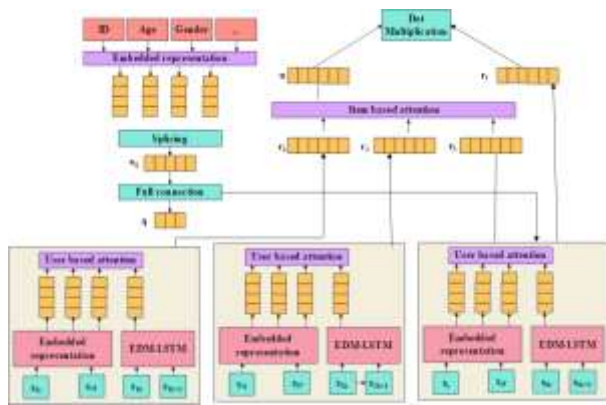


Figure 1: The framework of the proposed EDM-LSTM

The framework of the proposed EDM-LSTM is discussed as follows:

User attributes like ID, Age, Gender, etc., are encoded into an embedded representation. This process converts the categorical data into dense vectors which capture the important features of the user information.

The user-based attention processes the embedded user representation. It concentrates on the appropriate parts of the user's listening history by applying weights to various parts of the sequence based on their significance. The item-based attention mechanism is applied to music by estimating the significance of every track in a user's listening history to predict the next song the user is likely to prefer.

The dot multiplication step integrates the output from the user- and item-based attention mechanisms. This supports assigning the user's preferences to the music characteristics.

The music in the user's listening history is converted into embedded representations, which capture the features of every piece of music.

The EDM-LSTM module processes the embedded song representations. It applies the EDM to prioritize previous interactions, capturing the user's present preferences effectively.

The results from various EDM-LSTM modules are merged and passed through a fully connected layer. This step combines the data from various user-item interactions to design a comprehensive understanding of the user's preferences.

The combined representation (q) is further distinguished to produce the final recommendation score ($r_1, r_2, \dots, r_L$) for various music. This score represents the relevance of every music based on the user's listening history. Finally, the songs with the highest recommendation scores are suggested to the user.

Dataset

In this research, two standard music recommendation datasets are considered such as Lastfm [21], and 30Music [22]. Table 1 represents the statistics of the two utilized

datasets. The detailed description of these datasets is discussed in the following.

LastFM

The LastFM is the standard data which is obtained through the LastFM Application Programming Interface (API). Every sample in this data involves various attributes like timestamp, artist_name, user_id, artist_id, track_name (item_name) and track_id (item_id). The most three significant attributes like user_id, timestamp, and item_id are only taken for further processing because these enable the proposed method to capture the exact sequence and timing of user interactions. The session identity (session_id) is not applicable due to physical partitioning, which segments each user's interaction sequence into sessions lasting 30 minutes. The items with a duration of more than 5 minutes are filtered out and the users who have engaged with the platform less than 3 times are also filtered out.

30Music

30Music dataset is a real-time widely accessible standard collection of playlists and listening data. This dataset is acquired from the Last.fm API to recover the data from the internet radio stations. An original 30Music dataset involves 5 various attributes like session_id, song_id, user_id, playtime and ts. In this research, an attribute or column like playtime is removed. Each user has different session IDs, which are organized sequentially. Each event or item is integrated with a particular user and the similar session ID is integrated through using the temporal order for examination. Hence, various session sequences are created with lengths ranging from 30 to 199, and the average session sequence length is approximately 8.98 items.

Table 1. Statistics of the datasets

Attributes	Dataset	
	Lastfm	30Music
Number of users	277	4106341
Number of sessions	23230	2764474
Number of items	122816	210633
Number of events	683907	31351954
Minimum session per length	3	1
Average session per length	29.47	11
Maximum session per length	3997	198
Minimum session per user	3	1
Average per user	199	38
Maximum session per user	83.91	3.70

Recommendation based EDM-LSTM

In this research, the quality of the music recommendation is enhanced through developing the recommendation approach based on the EDM-LSTM. This research involves three parts such as music coding, user coding and recommendation. Initially, design a user model to learn music representations. Afterwards, designing the music model to learn user- preference representation according to the representation of the music listened to by the users. At last, the learned user and music representations are used to predict the probability of a user's preferences for the candidate music, enabling for

the recommendation of the top N tracks with more ratings to the users.

Music Coding

Presently, the coding approach of the music generally utilizes the embedded representation method to integrate an exclusive representation of the users listening to the music into condensed semantic feature vectors and utilize them as coding music features. This coarse-grained coding approach does not represent the fine-grained features of the music, which tends to ineffective music recommendations. User's preference for the song is generally affected by the multidomain features like music genre, singer and melody of the music. Furthermore, various users lead to perform various attention to the preferences of the similar music. Hence, it is important to differentiate user's various preferences for the music features from a fine-grained level. Thus, the music coding in an impairment approach in a user-based attention network.

User Coding

The user coding aims to learn the user's preference representation from the representation of music features listened to by the user. Various approaches such as weighted pooling, attention mechanism and Cyclic Neural Network (CNN) have been utilized to learn user preference. Here, the user attention model is utilized to learn user coding for tailoring the music recommendations to individual user preferences. Particularly, the EDM-LSTM is utilized to learn fine-grained mastery requirements among the user listening to the music. Concerning attaining the mastery requirement of user's listening behaviour, every position L in the user's listening order is determined, and described through c_l and each position codes of listening order are depicted through the matrix $c = [c_1, c_2, \dots, c_L]$, $C \in \mathbb{N}^{jj \times L}$. Hence, the music representation with the position-coding in the user's listening order is depicted as $F = C \oplus R$, where, $\oplus$ represents the element level addition operation $F \in \mathbb{N}^{jj \times L}$ respectively. User-based attention considers the individual user preferences through weighting interactions according to their significance, improving personalization. Item-based Attention focuses on the music items to match them with user preferences and determine similar items. As a result, these mechanisms significantly improve the quality and effectiveness of recommendations, resulting in a more engaging and personalized user experience.

Assuming the music feature representation f in the sequence as the query Q , key K and value in a self-attention mechanism. This mechanism captures contextual relationships within a user's interaction history, allowing the proposed method more refinement and context-aware recommendations. It is passed through the Fully Connected (FC) network for non-linear transformation. This approach establishes a relationship weight matrix that captures patterns in music listening behaviour, as formulated in Eqs. (1) to (3) as follows:

$$A = \text{Softmax}\left(\frac{QK^T}{\sqrt{N_{jj}}}\right) \quad (1)$$

$$QQ = \eta(\overline{F^T W^0} + B_q) \quad (2)$$

$$K = \eta(F^T W^K + B_k) \quad (3)$$

Where, $\eta(FW)$ represents the nonlinear activation function. Here, the ReLU function is utilized to enhance a nonlinear capability, $W_{QQ} \in \mathbb{N}^{i \times N_i}$, $W^K \in \mathbb{N}^{i \times N_i}$, $B_q \in \mathbb{L} \times N_i$, $B_k \in \mathbb{L} \times N_i$ and denotes the model parameter. $\sqrt{N_{jj}}$ scale represents a dot multiplication task. The results of the attention network are a matrix of $L \times N_i$ which is formulated in Eq. (4) as follows:

$$O = AF^T \quad (4)$$

To obtain the user's global preference representation, this research segments behaviour over time and then uses a fully connected (FC) network to learn the user's comprehensive preferences, as formulated in Eq. (5) as follows:

$$u_i = \text{concat}(O_1, O_2, \dots, O_L)W^u \quad (5)$$

Where, concat represents the vector splicing function, $W^u \in \mathbb{L} \times N_i$. The concat is a function that integrates the vectors end-to-end along a particular axis. The concatenation result is multiplied through the weight matrix W^u to express the user's comprehensive preference representation u_i respectively. Music coding offers the comprehensive features of every track, while user coding captures user interaction patterns and preferences. An integration of these coding allows the recommendation system to make accurate, personalized music recommendations through understanding both the user's details of music and behaviour. Recommendation systems propose music to users based on their past interactions and preferences, directing them to provide personalized choices. Then, the recommended results are provided for the prediction process.

Recommendation using Exponential Decay Moment based Long Short-Term Memory

The overall coding process of the recommendation system involves inputting user data to predict their future preferences based on past behavior. As compared to other Neural Networks, LSTM is effective at learning patterns in sequences, that makes them better understanding in which music is played. This contextual understanding leads to more accurate and personalized recommendations. LSTM is a kind of Recurrent Neural Network (RNN), it can record the long-term and short-term data simultaneously. The LSTM architecture involves various gates such as input gate, forget gate and output gate. It manages data by utilizing these gates to control information flow: an input gate incorporates new data, the forget gate removes irrelevant data, and an output gate updates the cell state. These operations allow the LSTM to capture and remember short-term and long-term dependencies in data. The relevant in-built operations comprise sigmoid and Tanh activations which

control and transform data within the network. These gates are used to retain essential long-term data while discarding unnecessary data. The LSTM gates are capable of regulating the data and forwarding it into the further unit. An input gate handles the new data to extend a further cell state and it performs in binary portions. An initial portion is a sigmoid layer, which controls an output value stored in a cell state. Another part is a Tanh layer which designs the new feature vector values stored in a cell state. A forget gate either completely discards data or retains it based on its values. An output gate updates a cell state data. With the structure of these gates, the data operates effectively by leveraging historical information to update the cell state. LSTM assumes prior historical values, and examines the current unfamiliar patterns by modifying itself based on the comprehensive patterns. Fig. 2 depicts the functionality of LSTM.

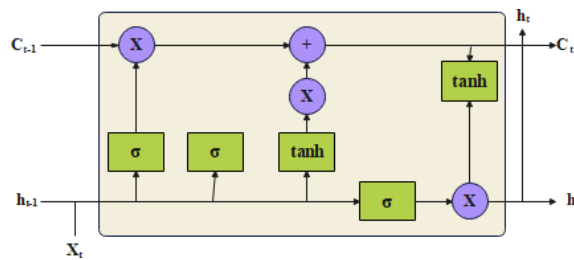


Figure. 2 Functionality of LSTM

In Fig. 2, X_t represents the input sequence, which is the system state monitoring data. h_t represents a result of a hidden layer, which is a learning outcome of every LSTM unit. f_t denotes a forgetting gate, i_t represents a input gate, O_t denotes a output gate, σ depicts a sigmoid function, $\tanh$ represents a activation function, W denotes a weight matrix.

Forget gate: LSTM performs the time-series data in a sequence manner. A forgetting gate function is to decide which has to be recollected and which has to be disregarded. An output of the forget gate is formulated in Eq. (6) as follows:

$$f_t = \sigma(W_f \cdot [h_{t-1}, X_t] + b_f) \quad (6)$$

Where, W_f denotes the weight of forget gate; b_f denotes a bias of forget gate.

Input Gate: The input gate function is to identify which parameters require to be updated and how to update. The result of an input gate is formulated in Eqs. (7) to (9) as follows:

$$i_t = \sigma(W_i \cdot [h_{t-1}, x_t] + b_i) \quad (7)$$

$$C_t = \tanh(W_c \cdot [h_{t-1}, x_t] + b_c) \quad (8)$$

$$C_t = f_t * C_{t-1} + i_t * C_t \quad (9)$$

Where, W_i , W_c depicts the respective weights of the input state and cell state. b_i , b_c denotes the bias; and C_t denotes a current cell state value.

Output Gate: The data obtains an output gate after being screened through a forget and input gate. An output gate function is to decide which data has to output. The result of an output gate is formulated in Eqs. (10) and (11) as follows:

$$O_t = \sigma(W_{OO} \cdot [h_{t-1}, x_t] + b_{OO}) \quad (10)$$

$$h_t = O_t \cdot \tanh(C_t) \quad (11)$$

where, W_{OO} denotes a weight of an output gate

b_{OO} denotes a bias of an output gate and h_t represents an output value of the current unit

This research proposes the improved optimization function called EDM with the LSTM approach, which aims to minimize the error results and improve the model performance. The EDM is a technique utilized in optimization approaches, particularly in the context of LSTM networks. The EDM involves adjusting the gradients during the training process by applying the exponentially decreasing weights to previous gradient updates. This supports prioritizing current information, thus enhancing the training efficiency and convergence of LSTM models through solving problems like vanishing or exploding gradients. In this proposed method, the optimization function performs after the estimation of the error and the propagation will begin that update the weights of matrices and cell states of LSTM. An Adam optimization function minimizes an error and achieves faster convergence through the iterative process of backpropagation. EDM produces a regularization effect by dynamically adjusting the learning rate, which prevents the LSTM from overfitting to the training data. It calculates gradients that denote the direction and magnitude of adjustments to model parameters, enabling efficient error reduction during training iterations. From this overall analysis, this section represents the process of improving music recommendations through encoding music and user features by proposing the EDM-LSTM framework. This prediction supports in recommending the top K music items that are most chosen by the user. Examining the user preference with the help of EDM-LSTM effectively enhances the recommendation performance.

Experimental Results

In this section, the effectiveness of the proposed EDM-LSTM approach is executed on Python 3.8 environment and the system configuration with Windows 10 OS, Intel Core i7 processor, 16GB RAM. The proposed method is validated through utilizing various assessment metrics named precision, Recall @ k, and Mean Reciprocal Rank (MRR) @ k. The mathematical formula of every metric is described in the following Eqs. (12) and (13).

Performance Analysis

In this section, the achievement of the proposed EDM-LSTM approach is estimated by using various assessment metrics by using two different music datasets Lastfm and 30Music. Tables 3 and 4 represent the performance analysis of the proposed method using Lastfm and 30Music dataset. The @10 and @20 represent the most relevant items in top 10 and 20

recommendations. Table 2. Hyperparameters of the EDM-LSTM

$$Recall = \frac{TP}{TP+FN} \quad (12)$$

$$MRR = \frac{1}{|QQ|} \sum_{i=1}^{|QQ|} \frac{1}{Rank K_i} \quad (13)$$

where, TP - True Positive; TN – True Negative; FP – False positive; FN – False Negative;

|QQ| represents the number of queries; *rank_i* denotes the rank position of the first relevant result for the *i*th query. *k* represents the top k recommendations. Table 2 provides the hyperparameter results of the EDM-LSTM.

Table 2. Hyperparameters of the EDM-LSTM

Hyperparameter	Values
Latent vector size	100 for each item
Kernel size	3 × 3
Batch size	100
Epoch	25
Activation function	SoftMax
Optimizer	Adam
Loss function	Categorical loss function
Learning rate	0.001

In this research, the existing DL approaches such as CNN, GRU, RNN and LSTM are compared with the proposed method to validate the results.

As compared to these existing methods, the proposed EDM-LSTM allows the recommender to significantly capture long-term dependencies in the sequential data of user listening behaviors, leading to more effective results in music recommendation. In Table 3, the achievement of the proposed EDM-LSTM method is presented with the different performance metrics on Lastfm dataset.

In Table 4, the achievement of the proposed EDM-LSTM method is presented with the different performance metrics on 30Music dataset.

Table 4. Analysis of proposed EDM-LSTM using 30Music dataset

Methods	Recall@10	MRR@10	Recall@20	MRR@20
CNN	25.843	21.305	31.934	21.359
GRU	27.896	23.361	33.202	23.337
RNN	29.412	24.494	34.943	25.595
LSTM	31.485	25.753	35.569	26.968
Proposed EDM-LSTM	33.303	26.300	37.384	28.393

Achievement of the DL-based existing recommendation approaches like CNN, GRU, RNN and LSTM are compared and tested with the proposed EDM-LSTM approach. The proposed EDM-LSTM approach attains better results with Recall@10 of 33.303, MRR@10 of 26.300, Recall@20 of 37.384 and MRR@20 of 28.393 respectively. However, the traditional LSTM attains the Recall@10 of 31.485, MRR@10 of 25.753, Recall@20 of 35.569 and MRR@20 of 26.968 respectively. These results show that the proposed EDM-LSTM approach attains better results as compared to the existing methods.

In Table 5, the analysis of the different techniques with LSTM using the 30Music dataset is presented. The different mechanisms like Step Decay (SD)- LSTM, Polynomial Decay (PD) -LSTM, Linear Decay (LD) -LSTM and Exponential Moving Average (EMA) are estimated with the proposed EDM-LSTM approach. As compared to other methods, EDM supports maintaining the temporal dynamics in music recommendation through smoothing the influence of previous data. The proposed EDM-LSTM approach attains better results with Recall@10 of 33.303, MRR@10 of 26.300, Recall@20 of 37.384 and MRR@20 of 28.393 respectively.

Table 5. Analysis of different techniques with LSTM using 30Music dataset

Methods	Recall@10	MRR@10	Recall@20	MRR@20
SD-LSTM	26.393	22.912	33.184	22.103
PD-LSTM	28.134	22.293	34.953	24.404
LD-LSTM	30.573	23.294	35.392	26.572
EMA -LSTM	32.858	24.049	36.303	27.758
Proposed EDM-LSTM	33.303	26.300	37.384	28.393

Figures. 3 and 4 represent the K-fold values of the proposed EDM-LSTM approach using Lastfm and 30Music datasets. The different K-fold values such as 3, 4, 5, 7 and 8 are considered to estimate the effectiveness of the proposed EDM-LSTM approach. As compared to the other K-fold values, the K=5 attains the better Recall@10 of 23.370 and 33.303 for both Lastfm and 30Music datasets due to the balanced trade-off between the bias and variance. With K=5, each fold contains 80% of the data for training and 20% for validation.

This distribution tends to reduce the variability in model performance estimates compared to lower K values, where each fold contains the minimum data for training or validation.

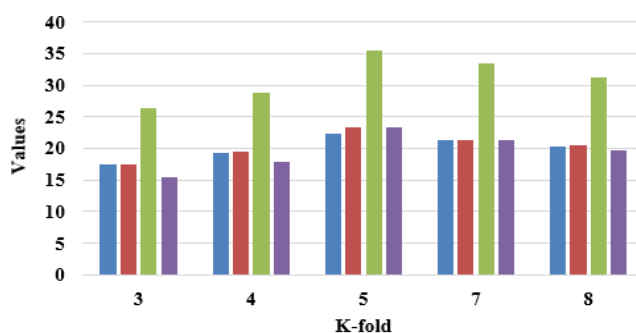


Fig Recall@10 MRR@10 Recall@20 MRR@20 Lastfm dataset

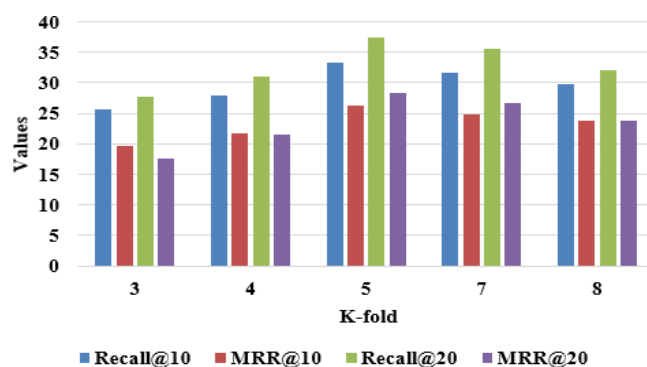


Figure. 4 K-fold analysis of the proposed EDM-LSTM using the 30Music dataset

Comparative Analysis

Table 6 compares the proposed EDM-LSTM approach with the existing approaches using different evaluation indices based on the three different datasets such as Lastfm, and 30Music. The existing methods such as SF-SKNN [16], AIRE [17], IR-MC [18] and GASM [19] are compared with the proposed EDM-LSTM by utilizing Recall@10, MRR@10, Recall@20 and MRR@20 respectively. The proposed EDM-LSTM performs the detailed music and user feature encoding, exploiting fine-grained representations for better performance in recommendations.

Table 6. Comparative Analysis

Dataset	Methods	Recall@10	MRR@10	Recall@20	MRR@20
LastFM	SF-SKNN [16]	0.2305	0.1712	NA	NA
	AIRE [17]	0.5682	NA	NA	NA
	IR-MC [18]	NA	NA	NA	0.346
	GASM [19]	21.818	21.582	33.320	21.829
	Proposed EDM-LSTM	22.370	23.303	35.394	23.315
30Music	SF-SKNN [16]	0.2787	0.1165	NA	NA
	IR-MC [18]	NA	NA	NA	0.346
	GASM [19]	32.733	25.901	36.096	26.136
	Proposed EDM-LSTM	33.303	26.300	37.384	28.393

Discussion

In this section, the achievement of the proposed EDM-LSTM approach and the drawbacks of the existing approaches are demonstrated. The limitations of the existing methods are as follows: the multiple ML approach [16] was overly complex and struggled with sparse user-item interaction matrices due to missing and non-zero entries, leading to unreliable similarity measures. The AIRE [17] heavily depends on the input data quality, inaccurate or noisy data leads to unreliable recommendations. The IR-based approach [18] is not prepared to maintain the changes in the relevance of items, impacting the time-awareness capabilities, so it does not inherently focus on the temporal dynamics of user preference. The GASM [19] was vulnerable to overfitting due to the complex architecture of the model, so it led to poor performance data. The MAR[20] model required large volumes of high-quality data to learn meaningful attention patterns. Hence, this research proposes the EDM-LSTM for the recommendation of music. The EDM-LSTM dynamically adjust to changes in user preferences over time, making the recommendation system more responsive to recent listening trends and enhancing user satisfaction.

Conclusion

Large music data and various listening behaviors present significant challenges for traditional methods in user-personalized recommendation situations. This research aims to propose a novel EDM-LSTM approach for the recommendation of music. EDM allows the recommendation system to prioritize recent user interactions and preferences, thereby enhancing the personalization of suggestions. LSTM with its gating mechanisms, effectively captures and retains long-term user interactions, enhancing personalized recommendations. This capability helps in understanding user behavior and adapting to changing preferences over time. By integrating EDM with the LSTM approach, the system effectively controls the sparsity of user listening data by enhancing an exponential decay mechanism to prioritize the most relevant historical interactions. This research involves three parts music coding, user coding and prediction. The experimental results show that the proposed EDM-

LSTM achieves better Recall@10 of 22.370 and 33.303 on Lastfm and 30Music datasets when compared to the existing methods. The future work will explore various deep learning approaches for music recommendation to improve overall performance.

Conflicts of Interest

The authors declare no conflict of interest.

Author Contributions

The paper conceptualization, methodology, software, validation, formal analysis, investigation, resources, data curation, writing-original draft preparation, manuscript - review and editing, visualization, and formal analysis have been done by 1st author. The supervision and project administration, have been done by 2nd author.

REFERENCES

- 1) Z. Liu, W. Xu, W. Zhang, and Q. Jiang, “An emotion-based personalized music recommendation framework for emotion improvement”, *Information Processing & Management*, Vol. 60, No. 3, p. 103256, 2023.
- 2) R.J.K. Almahmood, and A. Tekerek, “Issues and solutions in deep learning-enabled recommendation systems within the e-commerce field”, *Applied Sciences*, Vol. 12, No. 21, p. 11256, 2022.
- 3) R. De Prisco, A. Guarino, D. Malandrino, and R. Zaccagnino, “Induced emotion-based music recommendation through reinforcement learning”, *Applied Sciences*, Vol. 12, No. 21, p. 11209, 2022.
- 4) J. Li, Z. He, Y. Cui, C. Wang, C. Chen, C. Yu, M. Zhang, Y. Liu, and S. Ma, “Towards ubiquitous personalized music recommendation with smart bracelets”, *Proceedings of the ACM on Interactive, Mobile, Wearable and Ubiquitous Technologies*, Vol. 6, No. 3, pp. 1-34, 2022.
- 5) Z. Volovikova, P. Kuderov, and A.I. Panov, “Interpreting Decision Process in Offline Reinforcement Learning for Interactive Recommendation Systems”, In: *International Conference on Neural Information Processing*, pp. 270-286, 2023. Singapore: Springer Nature Singapore.

- 6) Y. Zhang, “Music recommendation system and recommendation model based on convolutional neural network”, *Mobile Information Systems*, Vol. 2022, No. 1, p. 3387598, 2022.
 - 7) S. Sahu, R. Kumar, M.S. Pathan, J. Shafi, Y. Kumar, and M.F. Ijaz, “Movie popularity and target audience prediction using the content-based recommender system”, *IEEE Access*, Vol. 10, pp. 42044-42060, 2022.
 - 8) O. Ghosh, R. Sonkusare, S. Kulkarni, and S. Laddha, “Music Recommendation System based on Emotion Detection using Image Processing and Deep Networks”, In: 2022 2nd International Conference on Intelligent Technologies (CONIT), pp. 1-5, 2022. IEEE.
 - 9) E. Sarin, S. Vashishtha, and S. Kaur, “SentiSpotMusic: a music recommendation system based on sentiment analysis. In 2021 4th International Conference on Recent Trends in Computer Science and Technology (ICRTCST), pp. 373-378, 2022. IEEE.
 - 10) R. Borges, and K. Stefanidis, “Feature-blind fairness in collaborative filtering recommender systems”, *Knowledge and Information Systems*, Vol. 64, No. 4, pp. 943-962, 2022.
 - 11) M. Moscati, E. Parada-Cabaleiro, Y. Deldjoo, E. Zangerle, and M. Schedl, “Music4All-Onion--A Large-Scale Multi-faceted Content-Centric Music Recommendation Dataset”, In: *Proceedings of the 31st ACM International*
 - 12) [S.K. Prabhakar, and S.W. Lee, “Holistic approaches to music genre classification using efficient transfer and deep learning techniques”, *Expert Systems with Applications*, Vol. 211, p. 118636, 2023.
 - 13) Z.Z. Darban, and M.H. Valipour, “GHRS: Graph-based hybrid recommendation system with application to movie recommendation”, *Expert Systems with Applications*, Vol. 200, p. 116850, 2022.
 - 14) B. Bakariya, A. Singh, H. Singh, P. Raju, R. Rajpoot, and K.K. Mohbey, “Facial emotion recognition and music recommendation system using CNN-based deep learning techniques”, *Evolving Systems*, Vol. 15, No. 2, pp. 641-658, 2024.
 - 15) M. Kleć, A. Wiczorkowska, K. Szklanny, and W. Strus, “Beyond the Big Five personality traits for music recommendation systems”, *EURASIP Journal on Audio, Speech, and Music Processing*, Vol. 2023, No. 1, p. 4, 2023.
 - 16) C. Kumar, and M. Kumar, “User session interaction-based recommendation system using various machine learning techniques”, *Multimedia Tools and Applications*, Vol. 82, No. 14, pp. 21279-21309, 2023.
 - 17) Y. Tao, C. Wang, and A.W. C. Liew, “Dynamic weighted ensemble learning for sequential recommendation systems: The AIRE model”, *Future Generation Computer Systems*, Vol. 149, pp. 162-170, 2023.
 - 18) A. Tofani, R. Borges, and M. Queiroz, “Dynamic session-based music recommendation using information retrieval techniques”, *User Modeling and User-Adapted Interaction*, Vol. 32, No. 4, pp. 575-609, 2022.
 - 19) H. Weng, J. Chen, D. Wang, X. Zhang, and D. Yu, “Graph-based attentive sequential model with metadata for music recommendation”, *IEEE Access*, Vol. 10, pp. 108226-108240, 2022.
 - 20) W. Lu, “Design of a music recommendation model on the basis of multilayer attention representation”, *Scientific Programming*, Vol. 2022, No. 1, p. 7763726, 2022.
- Lastfm dataset link:
<https://www.kaggle.com/datasets/hoandan/lastfm> m.
 (Accessed on 09/07/2024). 30Music dataset
 link: <https://recsys.deib.polimi.it/datasets/>
 (Accessed on 09/07/2024).