

Leveraging Social Trust Matrix Factorization for Enhanced Recommender Systems

Rosni Lumbantoruan, Petrus Sinaga, Ignatia Hutagalung, Christy Sihombing, Humasak Simanjuntak, Arnaldo Marulitua Sinaga

Abstract—The growing of digital industry such as e-commerce has raised the issue of users facing the challenge in determining their needs due to the increasing number of available products options. We believe that the user's network or user's neighborhood can help to alleviate this issue by recommending items that the closest neighbors have purchased, yet we argue that the indirect influence of trust should also be considered. In this study, we adopt an approach that uses trust matrix factorization and collaborative filtering models to identify the influence of trust amongst users. This trust data will be mined to identify pattern in which a target user is influenced by other users that are in her networks. Trust matrix factorization is used to model the level of trust between users based on existing social relationships. Meanwhile, the collaborative filtering model is used to identify patterns of similarity in consumer characteristics. Through this approach, we aim to see which social recommendations are better suited to represent user preferences. Here, we proved that the proposed model offers more personalized and relevant recommendations by considering the indirect influence of trust on top of trust matrix factorization and collaborative filtering models. The results of this study have important implications for the development of effective and accurate recommendation systems. By highlighting the role of trust influence, this research provides valuable insights into understanding user interactions and designing better recommendation systems in the future.

Index Terms—Recommender system, Data mining, Trust user popular, Trust syllogism, Knowledge discovery, TrustMF, Social recommendation.

I. INTRODUCTION

THE recommendation system plays a crucial role in the decision-making process and is commonly employed by online platforms such as Amazon, Netflix, and Youtube. The concept of this recommendation system is an information filtering technique used to predict user preferences through personalization based on their past experiences. In essence, the

recommendation system uses the knowledge of a user's previous interactions with the platform to suggest items, products, or content that are likely to align with their individual and interests. By analyzing user behavior and preferences, the system can make personalized recommendations, thereby enhancing the user experience and increasing engagement on the platform [1]. In the commonly used recommendation systems (RS), there are three major categories: content-based (CB), collaborative filtering (CF), and hybrid methods. Among these approaches, Using the collaborative filtering technique is the most popular. The reason for its popularity is that this approach relies on past rating information to build models without requiring access to external data. So, collaborative filtering leverages historical user-item interactions to make personalized recommendations. It does not rely on external information and benefits from the idea that users with similar tastes will have similar preferences for items. This makes it a widely used and effective method in recommendation systems. [2]. Despite being widely used, collaborative filtering approach has a drawback in providing recommendations due to sparse data in the rating dataset, which affects the system's predictive performance. When the available rating data is limited and sparse, the system faces challenges in accurately identifying user patterns and preferences [3]. To overcome the issue of data sparsity, a collaborative approach using matrix factorization (MF) is employed and some research that working on top of this MF has shown a better performance in recommender systems such as in [4, 5].

However, conventional MF solely focuses on describing the interaction patterns between items and users, without taking into account contextual information, such as the connections between users. This contextual information is vital as it reflects how users are linked and share similar preferences. For instance, when someone is watching a movie, they often seek and consider recommendations from their close friends or trusted acquaintances rather than solely relying on the recommendation system [6]. Many studies have focused on exploring user relationships using direct trust-based MF, where a user simply assigns a trust value to another user. However, in real life, a person's trust relationship can be influenced by various factors, including the people we trust and the influence of those trusted by the people we trust. This phenomenon can be explained using the principle of syllogism. In the principle of syllogism, the trust we place in someone can spread through a network of social relationships. For instance, if we trust our close friend, and that close friend trusts someone else, we are more likely to extend our trust to that person as well. On the

R. Lumbantoruan is with Faculty of Informatics and Electric Engineering at Institut Teknologi Del (IT Del), Indonesia (e-mail: rosni@del.ac.id*).

P. Sinaga, I. Hutagalung, and C. Sihombing were students of Faculty of Informatics and Electric Engineering, Department of Information System at Institut Teknologi Del (IT Del), Indonesia (e-mails: iss19025@students.del.ac.id, iss19043@students.del.ac.id, and iss19061@students.del.ac.id).

H. Simanjuntak is with Faculty of Informatics and Electric Engineering at Institut Teknologi Del (IT Del), Indonesia (e-mail: humasak@del.ac.id).

A. M. Sinaga is with Faculty of Vocational Studies at Institut Teknologi Del (IT Del), Indonesia (e-mail: aldo@del.ac.id).

other hand, if the person we trust is trusted by someone else, it can also impact our trust in that person. In essence, these trust relationships between individuals form a complex network that goes beyond direct connections between individuals [6]. In this research, we explore the utilization of trust relationships between users, both directly and indirectly, employing syllogisms to observe their influence on predictions within the recommender system. To compare the effects of these two types of trust, we employ the trust-based matrix factorization method (Trust MF). Trust MF is a model that leverages trust relationships to enhance the accuracy of predictions, comprising two components: the Truster model and the Trustee model. The Truster model utilizes rating data, and during prediction, it involves multiplying the representations of user features and item features. On the other hand, the Trustee model employs user trust data by multiplication between users to obtain predictions. The outcomes of both the Truster model and the Trustee model are combined to yield an improved recommendation prediction result. This research provides contributions on comparing the effect of direct trust and indirect trust using Trust MF and evaluate the Trust MF model for trust-based recommendation systems. that will be evaluated using the values of precision, recall, F1-score, MAE and RMSE.

II. RELATED WORK

A. Recommender Systems

Collaborative Filtering (CF) is a powerful prediction technique utilized in recommendation systems to provide personalized recommendations based on user opinions, interests, and preferences. By building a database of user ratings for various items and leveraging the similarity between users, CF can accurately identify relevant recommendations for individual users [4, 5, 7]. There are two main categories of CF: User-based collaborative filtering and Item-based collaborative filtering. In the former, recommendations are made based on either high ratings or similarity to other users, while the latter uses trained models to uncover patterns within the input data [8]. The underlying principle of CF revolves around gathering and processing substantial amounts of user-rated items with similar tastes. This data is then used to establish similar relationships among users, often referred to as neighbors. Consequently, CF algorithms excel at suggesting items or products to users based on the preferences of others who share similar interests [4, 9].

To address the challenges posed by incomplete utility matrices in recommendation systems, Matrix Factorization (MF) is one of the recommended techniques. It effectively breaks down a large matrix into smaller matrices and tackles the issue of sparse data resulting from many users only providing ratings for certain items. This incomplete data can significantly impact the accuracy of user-item recommendations based on their interests. However, with MF, a personalization model is employed to map users and items into latent factor variable dimensions for prediction [10]. Trust-based social networks play a crucial role in further improving the recommendation system. Users have social relationships in the form of a random

graph, and when trust is strong enough, they will always receive recommendations, even for items that are rarely recommended in the category [11].

CF and MF work hand-in-hand to provide powerful and tailored recommendations to users based on their opinions, interests, and preferences [10]. Meanwhile, CF leverages the similarity between users to identify relevant recommendations, MF helps fill in the gaps in the utility matrix, enhancing the accuracy and effectiveness of the recommendation system [10]. The combination of these techniques ensures that users receive personalized and relevant recommendations, contributing to an improved user experience in various fields such as social media, e-commerce, and more.

This model is considered with each agent assigning values to each item, with the range between $[-1,1]$. Negative values indicate dislike and positive values indicate liking towards the recommended item. In giving these values, consumers will provide ratings after trying the product, and the given ratings will be used to calculate the similarity. The calculation of similarity value between consumers and items is formulated in equation (1).

$$\omega_{i,j} = f(x) = a_0 + |f_{ai} - f_{aj}| \quad (1)$$

Users have social relationships in the form of a random graph when trust is strong enough, user will always get recommendations even if they have items that are rarely recommended in the category [11, 12]. This network uses a peer-to-peer trust determination algorithm to count trust members that are not connected in a peer-to-peer manner algorithm to calculate trust members that are not directly connected [11]. To complement the recommendation system further, trust-based social networks play a crucial role. Users have social relationships in the form of a random graph, and when trust is strong enough, they will always receive recommendations, even for items that are rarely recommended in the category [13]. Trust network's influence is integrated into the MF process, enabling the consideration of trust values in user-item recommendations [14, 15].

Trust MF calculation combines the rating matrix and the trust matrix which then performs the MF to obtain latent factors. In the rating matrix, the latent factors are represented in the form of U and V which refer to user preferences and item characteristics, respectively. In the trust matrix, the latent factors are represented as B and W , which refer to user preferences, respectively. The trust relationship between users is not always reciprocal or commonly referred to as asymmetric. Therefore, each user is mapped to two specific latent vectors which are the truster feature vector B_i and the trustee-specific feature vector W_i . B_i and W_i signify the “trusting others” and “trusted by others”.

Truster model reviews items of interest by users and trust relationships based on opinions that build trust bonds. Through these opinions, a user can be influenced by other trusted users. The “truster” model is used to characterize the influence of specific users’ opinions on other users. It can be observed in the dataset that the users involved in the rating matrix R and the

trust matrix T are the same, allowing them to share specific latent feature space during MF.

Trustee model aims to characterize the influence of user i 's opinions on others who trust user i when making decisions. By seeking the specific latent feature space of trusted users (trustees) W , through the predicted user. Research has shown that MF by incorporating trust, TrustMF, has successfully overcomes the data sparsity problem on the highly sparse dataset.

B. Data Mining

Many data mining techniques are incorporated in order to identify patterns or retrieve the user or item profile in recommender systems. The most common data mining techniques used for this purpose is clustering such as K-Means. In [4], K-Means clustering was used to identify users most similar in reviewing purchased items in terms of their used vocabulary or words. In this research, we will also use group users based on their popularity and association rule mining to further increase the user's trust network.

III. METHODOLOGY

A. Problem Definitions

Matrix factorization (MF) is commonly used to predict unknown ratings in a user-item rating matrix. This matrix consists of users (U) and items (I), with each cell representing the rating given by a user to an item, usually on a scale of 1 to 5. If a rating is missing, it means the user hasn't rated or purchased that item. MF maps users and items into a low-dimensional latent feature space, which represents their interests or preferences. This approach helps predict missing ratings by identifying latent features specific to users and items.

In addition to the user-item rating matrix, there's a trust network represented by the neighbor trust relationship matrix (T). This matrix contains trust values between users, ranging from 0 to 1, denoting the level of trust one user has in another. Trust relationships are asymmetric, meaning if user P trusts user Q , it doesn't necessarily imply that Q trusts P in the same way. The objective of this research is to incorporate user trust values into the prediction process and assess how much the trust relationships influence the recommendations. By incorporating data mining techniques, we will group users that are potentially will be trusted by the target user, we named this trust scenario as trust user popular. The other alternative is by employing rule mining technique using the syllogism to find the other users that might also influence the target users, we named this approach as trust syllogism.

B. Proposed Model

The general architecture of the model including the proposed trust syllogism and trust user popular is depicted in Fig. 1. Data from the Internet in terms of rating and trust data is the input for the proposed model. This research uses a MF algorithm coupled with a trust context or called Trust MF. Trust data will be used to mine the patterns in which a user is influenced by others, thus the model will consider recommending the connected users' items to the target user.

The TrustMF approach is a development approach of conventional MF on the background of its weakness, which only focuses on user preferences for items but does not consider the trust relationship that users have. In real and virtual life, trust between users has a very crucial role in the decision-making process. MF has the basis that trust information can improve relevant prediction results by providing additional information about user preferences. The trust relationship between users will be modeled in the form of a trust matrix and then be combined with a rating matrix.

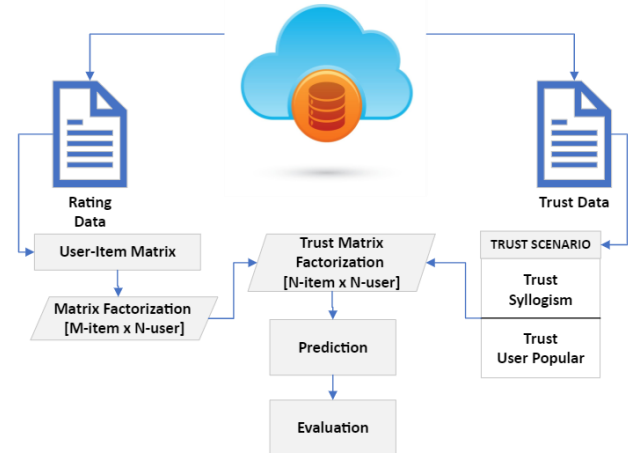


Fig. 1 The General architecture of the model

According to research [1] in making predictions, the Trust MF method consists of two parts, namely the truster model and the trustee model as a link between ratings and trust through mapping the latent space dimension (d) to the same size of each item and user. Before starting the training process, pre-processing of the trust mf model is carried out by collecting data where the data collected is rating data and user trust data. Furthermore, the rating data will be normalized using R_{Max} into a range between 0 and 1 in order to facilitate calculations in the training and prediction process. If it has gone through normalization, then the rating data will be divided into training sets and testing sets. Then the latent feature matrix B , V and W are initialized to train and test the truster model and trustee model. The truster model aims to predict ratings based on user characterization using matrix B as a latent space feature of users in assessing an item using matrix V through the influence of trust in other users using matrix W . Rating prediction in this truster model is based on multiplication or dot product of matrix B and matrix V using rating data to form a user-item matrix or in other words MF with a general form.

Meanwhile, the trustee model is a model that aims to predict ratings based on the characteristics of the influence of user trust using the latent features of the W matrix with other users, namely the B matrix in rating an item using the V matrix feature. Rank prediction in the trustee model by multiplying or product between matrix B and matrix W through trust data to form a trust matrix. During the model training process, we aim to find the best parameter value by using an objective function or loss function. This function helps measure and evaluate the Trust MF model's performance error rate. To minimize this

objective function, we employ the SGD optimization algorithm, which is iteratively applied until it approaches the minimum value. A lower value of the objective function indicates that the Trust MF model is more accurate in predicting user ratings for items that have never been rated before. The values obtained from the training process will be used to populate the matrix features in the truster and trustee models, respectively.

Based on previous research [2], the Trust MF model is used in conducting this research experiment which can be seen in Fig. 1. This research has differences with previous research on the use of datasets, which are obtained through two scenarios to see the effect of using trust either directly using the most trusted users (trust user popular) or further or indirect trust (trust syllogism).

Trust user popular, given the data from previous research [2], in this scenario, the data is added by identifying a list of users who receive the most trust. Thus, each user will have additional trust, not only to their current direct trust but also to the list of users who are trusted the most frequently. The reason for using popular users is to provide recommendations that are more personalized and have a big influence because they are a source of trust for many users so they can improve the results of trust predictions.

Trust syllogism, data obtained through the calculation of logical reasoning related to the trust relationship between users using baseline data. The purpose of forming this dataset is to find out the potential influence of trust from users who are in the trust chain with the assumption that it can further improve the results of recommendations. In real life, this syllogism trust can be illustrated as the following example. For example, Jane wants a haircut, then Jane entrusts her friend Alice to find a good hairdresser to cut her hair, so Alice recommends Mariah as a good hairdresser. Jane will automatically trust Mariah to cut her hair. The effect of the syllogism concept is that the number of transactions increases by 19 million and overcomes data scarcity in the baseline data by 98.82%.

IV. EXPERIMENTS EVALUATION

A. Experimental Evaluation

In this study we use the Epinions dataset which is widely used in research related to trust recommendation systems and is taken from the site http://www.trustlet.org/downloaded_epinions.html. This dataset contains product rating information and trust relationships between users. The rating dataset is one part of the Epinions dataset which contains the rating value given by users on items. This dataset consists of 40,163 users who gave ratings on 139,738 items. The rating distribution obtained is rating 5 is the maximum value given by users which is 301,053 and the minimum value is obtained by rating 1 of 43,228 which reflects that the rating distribution is not balanced as can be seen in Fig. 2.

Trust dataset is a dataset containing trust statements between one user and another. These users consist of users who provide or as a source of trust value (source) and users who receive the trust value (target). On average, each source user has 9.88 trust

relationships (target users), with a median of 2 and a standard

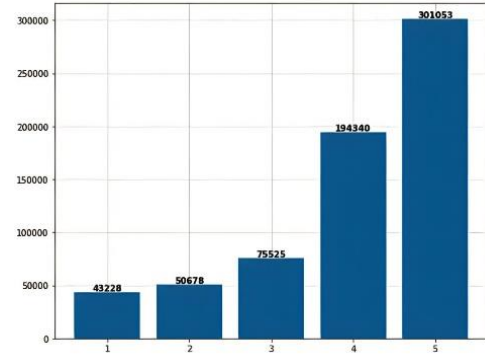


Fig. 2 The distribution of rating of Epinions dataset

deviation of 40.09. The scarcity rate of this data is 99.97% which falls into the high category. The syllogism is a dataset created through a derivative of the Epinions dataset using the concept of mathematical logic calculations to get a further trust relationship and determine the magnitude of the influence of trust in improving recommendations. In conducting the evaluation, we used the appropriate parameters to get the best results in the experiment.

B. Evaluation Methodology

In conducting the evaluation, we used the appropriate parameters to get the best results in the experiment. We set the same latent dimension space, $d = 10$, regulation parameter $\lambda = 0.001$, and $R_{\text{Max}} = 5$ for rating value normalization. Then, we divided the data equally in all our experiments where the training data is 80% and the rest is for testing data. The process uses a different number of epochs.

C. Experimental Results

Here, we assess the model's performance in terms of effectiveness compared to the baseline by using two evaluations namely to address 1) how relevant the model is in providing the recommendations to users in terms of Precision, Recall, and F1 score. Precision describes the percentage of items favored by users from all recommended items. Recall refers to the percentage of user-preferred items that appear in recommendations from all items preferred by users, while F1 is a combination of precision and recall evaluations; 2) we evaluate the prediction accuracy using Mean Average Error (MAE) and Root Mean Square Error (RMSE).

The effect of top-N popular users

The original trust dataset is supplemented with the results of user-context influence, focusing on users who receive the highest trust from others. Here, we assess the effect of the number of top users, Top-N User Popular scenario, to provide relevant recommendations based on user preferences and the preferences of users who are most trusted. The evaluation includes different top user variations, such as top 5, top 10, top 15, and top 20 users. As depicted in Fig. 3 and Table I, we can see that $N=15$ has the highest Precision, Recall, and F1-score. provides a better balance between recommendation quality

(precision) and the system's ability to find relevant items (recall).

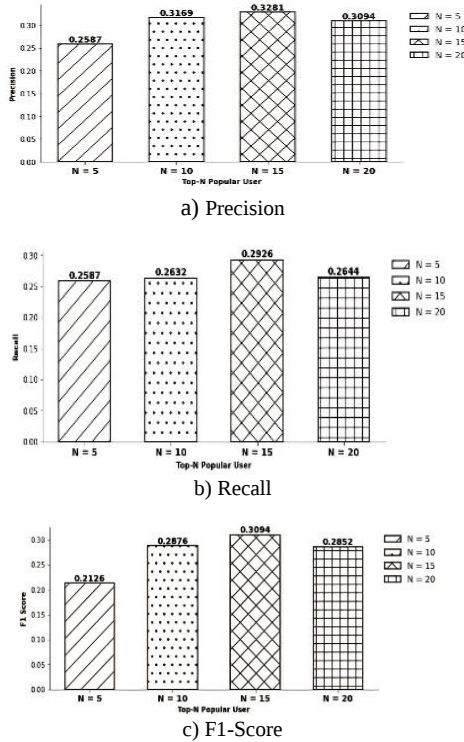


Fig. 3 Evaluation results of trust top-N popular user in terms of (a) Precision, (b) Recall, and (c) F1-Score.

TABLE I
RATING PREDICTION EVALUATION RESULTS

Dataset	Top-N	Metrics		
		Precision	Recall	F1-Score
Trust Original	5	0.2587	0.2587	0.2126
	10	0.3169	0.2632	0.2876
	15	0.3281	0.2926	0.3094
	20	0.3094	0.2644	0.2852

This indicates that adding 15 popular users as target user IDs. When N is increased to 20, the recommendation system seems to experience a decrease in performance due to the possibility of less relevant recommendations or neglecting some relevant items. This decrease in performance may occur because as N becomes larger, there is a possibility of adding popular users that significantly influence the MF, leading to suboptimal recommendations for the general users. Additionally, increasing the data in the system makes MF model more complex, which can result in overfitting or difficulty in finding meaningful patterns.

Therefore, N=15 is a better choice as it provides better performance than lower N values (N=5) and higher N values (N=20) in terms of precision, recall, and F1-score while maintaining a good balance between recommendation quality and the ability to find relevant items. When N=15 additional popular users are included as target user IDs, more relevant data is available to train the MF. This gives better opportunities for the User matrix (P) and Item matrix (Q) to understand user preferences and represent items more accurately in hidden vector form. However, when N=20 additional popular users are included, some negative effects may arise such as:

- Data noise: With a larger amount of data, there is a possibility of data noise or inaccurate data that can affect the quality of MF. This data noise can lead to poor results and affect the recommendation system's performance.
- Overfitting: A large amount of data can cause overfitting, where the MF "over-adapts" to the training data. This means the model may capture small details or noise in the training data but may not generalize well to unseen test data.
- Complex computations: As the data size increases, the MF process becomes more complex and time-consuming. This can slow down the recommendation system's performance and affect efficiency in providing recommendations to users.

Thus, referring to the experiment, we conclude that N=15 is the best setting to filter the number of popular users to employed.

We also assess this the effect of top popular users in terms of MAE and RMSE as can be seen in Table II and Fig. 4. Table II demonstrates that each application of top user variance has a distinct evaluation value. In the context of error evaluation, the lower the MAE value, the better the performance of the recommendation experiment. In the context of error evaluation, the lower the MAE value, the better the performance of the recommendation experiment, as well as the RMSE value. In the context of error evaluation, the lower the MAE value, the better the performance of the recommendation experiment, as well as the RMSE value, but this value is used to measure the average error between the predicted value and the actual value by taking the square root of the value. In both figures of the metric evaluation results error metric evaluation results, each has an x-axis that represents the overall comparison of the experiments from the top popular users and the y-axis represents the range of each value.

TABLE II
MAE AND RMSE ERROR EVALUATION RESULTS

Dataset	Top-N	Metrics	
		MAE	RMSE
Trust Original	5	1.1500	1.4425
	10	1.1296	1.4746
	15	1.1410	1.5163
	20	1.1412	1.4938

From Table II can be seen that top-N with a value of 15 still has the best value of the others. The following is displayed in the form of a bar chart of the evaluation results of MAE and RMSE.

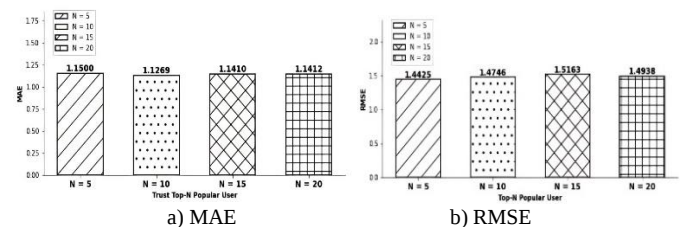


Fig. 4 Evaluation results of trust top-N popular user in terms of (a) MAE and (b) RMSE.

In both metric evaluation result figures errors, each of which has an x-axis representing the overall comparison of the experiments from the top popular users and the y-axis representing the range of values of each evaluation. The best

value in MAE is obtained in variation 10 which is 1.1269 whereas with other top user variations it has a difference range of 0.01-0.03. This means that the top 10 models are better at describing the trust relationship to produce predictions close to the original value. Whereas the best value in RMSE is obtained in variation 5, namely 1.4425 where the range of differences with other top users is 0.01-0.03. The range of differences with other top users is 0.03-0.07 which means it helps to reduce the influence of other users who may have a lower level of trust, lower level of trust.

Furthermore, the results and discussion of error evaluation metrics such as MAE (Mean Absolute Error) and RMSE (Root Mean Squared Error) on the trust-based recommendation model using TrustMF provide an overview of how close the model prediction results are to the actual rating value.

The comparison to the baselines

In this experiment, we compare two models we proposed in this research namely: 1) trust top user popular and, 2) trust syllogism with the baseline TrustMF in terms of Precision, Recall, F1 score, RMSE and MAE. The results of these experiments are depicted in Fig. 5. The description of each result figure for each evaluation metric is the x-axis. X-axis explains the comparison of each experiment consisting of baseline (original trust data), trust 15 top user popular and trust syllogism, y-axis is a scale score that reflects the evaluation results of the three experiments.

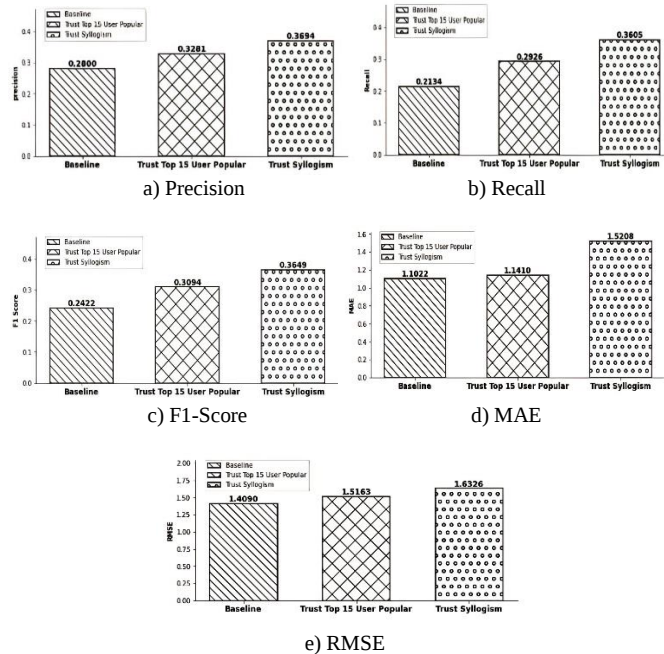


Fig. 5 Performance comparison of baseline trust MF vs. trust top 15 user popular vs. trust syllogism in terms of (a) Precision, (b) Recall, (c) F1-Score, (d) MAE, (e) RMSE.

It can be seen that the Trust Syllogism experiment provides the best results compared to the other two experiments, namely baseline and Trust 15 Top User Popular, based on the Precision, Recall, and F1 Score evaluation metrics with 0.3694, 0.3605, and 0.3649, respectively. Performance improvement from Trust

15 Top User Popular to Trust Syllogism can be explained by the use of Trust Syllogism. By utilizing the information between users, the MF model can discover more complex and accurate relationships between users and items. The concept of trust allows the recommender system to not only consider the user's direct preferences towards items, but also consider the preferences of people trusted by the user preferences of people whom the user trusts indirectly. Thus, the Trust Syllogism provides broader and more relevant information in assessing the user's preferences, which leads to more accurate recommendations that are more in line with the user's preferences and more in line with the user's preferences.

The results indicate that the use of Trust Syllogism has a positive impact on enhancing the accuracy of rating predictions, although no significant differences were observed. However, during the error evaluation using MAE and RMSE metrics, the best performing values were observed in the baseline, showing a difference range of 0.2 to 0.4 for trust syllogism. This observation could be attributed to several factors. Firstly, the baseline utilizes direct trust datasets, which are considered more reliable due to their extensive utilization in prior research on social recommendation systems. Conversely, the trust syllogism dataset is self-generated and yet to be validated. Additionally, the complexity of the trust syllogism data, involving more than two users in the trust chain and having a substantial size, seems to impact the model during training and consequently yields higher error values.

Despite these observations, the graphs substantiate that incorporating the trust syllogism concept into the Trust MF model can indeed enhance the accuracy of rating predictions. However, further research is imperative using alternative datasets to validate the research findings and explore different methodologies to assess the strengths and weaknesses of the model.

V. CONCLUSION

In this paper we propose Trust MF for use in providing more effective recommendations to users. Using the Trust MF methodology, we try two new methods, Trust Syllogism which is obtained by calculating the logical reasoning related to the trust relationship between users, namely trustee and truster. Meanwhile, the next methodology is Popular Users which is obtained from calculating the search list of users who receive the highest frequency of trust from the source user. And from the two methods that have been carried out, it can be concluded that the experimental results show that trust does not improve the MAE and RMSE metrics by producing large values. In other words, the recommendations provided are not accurate.

However, in terms of Precision and Recall, the model is better at recognizing and recommending relevant items to users. In contrast, using trusted original data results in smaller MAE and RMSE errors, thus providing more accurate recommendations. In addition, the experimental results show that trust does not improve the MAE and RMSE metrics by producing large values. In other words, the recommendations provided are not accurate. On the other hand, using the original trust data results in higher MAE and smaller RMSE errors, thus providing more accurate recommendations.

ACKNOWLEDGMENT

This research is fully supported by Institut Teknologi Del through *Lembaga Penelitian dan Pengabdian Masyarakat (LPPM)*.

REFERENCES

- [1] H. Liu, L. Jing, J. Yu, and M. K. Ng, "Social recommendation with learning personal and social latent factors," *IEEE Trans. Knowl. Data Eng.*, vol. 33, no. 7, pp. 2956–2970, 2021, doi: 10.1109/TKDE.2019.2961666.
- [2] B. Yang, Y. Lei, J. Liu, and W. Li, "Social collaborative filtering by trust," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 39, no. 8, pp. 1633–1647, 2017, doi: 10.1109/TPAMI.2016.2605085.
- [3] H. Ma, H. Yang, M. R. Lyu, and I. King, "Sorec: social recommendation using probabilistic matrix factorization," *In Proceedings of the 17th ACM conference on Information and knowledge management*, pp. 931–940, October 2008.
- [4] R. Lumbantoruan, P. Simanjuntak, I. Artonang, E. Simaremare, "TopC-CAMF: A top context-based matrix factorization recommender system," *Jurnal Nasional Teknik Elektro dan Teknologi Informasi*, 11(4), pp.258–266, 2022.
- [5] R. Lumbantoruan, X. Zhou, Y. Ren, and Z. Bao, "D-cars: A declarative context-aware recommender system," *In IEEE International Conference on Data Mining (ICDM)*, pp. 1152–1157, November 2018.
- [6] S. Xu, H. Zhuang, F. Sun, S. Wang, T. Wu, and J. Dong, "Recommendation algorithm of probabilistic matrix factorization based on directed trust," *Comput. Electr. Eng.*, vol. 93, no. March, p. 107206, 2021, doi: 10.1016/j.compeleceng.2021.107206.
- [7] K. Patel and H. B. Patel, "A state-of-the-art survey on recommendation system and prospective extensions," *Comput. Electron. Agric.*, vol. 178, no. June, p. 105779, 2020, doi: 10.1016/j.compag.2020.105779.
- [8] Y. Koren, S. Rendle, and R. Bell, "Advances in collaborative filtering," *Recommender systems handbook*, pp.91–142, 2021.
- [9] M. G. Ozsoy and F. Polat, "Trust based recommendation systems," *Proc. 2013 IEEE/ACM Int. Conf. Adv. Soc. Networks Anal. Mining, ASONAM 2013*, pp. 1267–1274, 2013, doi: 10.1145/2492517.2500276.
- [10] A. K. Sahoo, C. Pradhan, R. K. Barik, and H. Dubey, "DeepReco: Deep learning based health recommender system using collaborative filtering," *Computation*, vol. 7, no. 2, 2019, doi: 10.3390/computation7020025.
- [11] Chae, Dong-Kyu, Jihoo Kim, Duen Horng Chau, and Sang-Wook Kim, "AR-CF: Augmenting virtual users and items in collaborative filtering for addressing cold-start problems," *In Proceedings of the 43rd International ACM SIGIR Conference on Research and Development in Information Retrieval*, pp. 1251–1260. 2020.
- [12] R. Lumbantoruan, X. Zhou, and Y. Ren, "Declarative user-item profiling based context-aware recommendation," *In International Conference on Advanced Data Mining and Applications*, pp. 413–427, November 2020.
- [13] M. Jamali and M. Ester, "A matrix factorization technique with trust propagation for recommendation in social networks," *RecSys'10 - Proc. 4th ACM Conf. Recomm. Syst.*, pp. 135–142, 2010, doi: 10.1145/1864708.1864736.
- [14] A. S. Tewari, J. P. Singh, and A. G. Barman, "Generating top-N items recommendation set using collaborative, content based filtering and rating variance," *Procedia Comput. Sci.*, vol. 132, no. Iccids, pp. 1678–1684, 2018, doi: 10.1016/j.procs.2018.05.139.
- [15] Anelli, Vito Walter, Alejandro Bellogín, Tommaso Di Noia, Dietmar Jannach, and Claudio Pomo, "Top-n recommendation algorithms: A quest for the state-of-the-art," *In Proceedings of the 30th ACM Conference on User Modeling, Adaptation and Personalization*, pp. 121–131. 2022.