

Utilizing Textual Reviews for Visualizing and Understanding User Preferences

Dang Pham

Department of Computer Science
New Mexico State University, Las Cruces, USA
dangpnh@nmsu.edu

Tuan M. V. Le

Department of Computer Science
New Mexico State University, Las Cruces, USA
tuanle@nmsu.edu

Abstract—Latent factor models are widely used in recommender systems. In these models, users and items are represented as vectors in a joint latent factor space. The inner products of user vectors and item vectors are used to model the user-item interactions (e.g., ratings). A review is often posted by the user to explain the given rating. Therefore, reviews can be used to understand how users rate the items and to interpret the latent dimensions of user and item vectors. In this paper, we propose a probabilistic model that learns latent vectors of users and items in a two- or three-dimensional space for visualization. Our proposed model also extracts review topics and visualizes them in the same visualization space for interpreting the ratings. We model the user-item interactions by using the distances between users and items in the visualization space. Extensive experiments using several real-world datasets demonstrate the effectiveness of our proposed model in recommendation and visualization tasks.

Index Terms—latent factor models, visualization, topic models

I. INTRODUCTION

In recommender systems, users often express their preferences by giving ratings and writing textual reviews to explain the ratings they give. To understand user preferences, latent factor models are widely used to analyze the ratings [1], [2]. In these models, users and items are represented as vectors in a latent factor space. The ratings are modeled as the inner products of user and item vectors. Given training data, the latent vectors are typically learned by minimizing the mean square error. While the learned latent vectors are good for rating prediction, one shortcoming is that we may not have a clear understanding of why users would give such ratings. In this regard, reviews can be a great source for extracting discussed topics and aspects that could provide an explanation on the given ratings. Therefore, there have been several works that leverage reviews for understanding rating factors and improving the performance of recommender systems [3]–[6].

For making sense of how users rate items, visualization is an effective means to present user preferences and explore item recommendations. Visualization approach has shown its effectiveness in different tasks such as presenting and manipulating user preferences [7], exploring social recommendations [8], visualizing explanations in recommender systems [9]. In these methods, one possible way to display users and items is by embedding their latent vectors to a 2- or 3-dimensional space using a dimensionality reduction method such as Multidimensional Scaling [7]. However, most of the existing visualization methods ignore the reviews when performing the dimensionality reduction of user and item vectors. Therefore, the explanations based on reviews may not be transferred to and visualized in the visualization space for making sense of the item recommendations. Although document visualization methods can be utilized to visualize topics in reviews [10], [11], they do not model users, items, and ratings. Therefore, the output visualization may be not relevant and reflect well the user rating behavior.

To address the above limitations of the existing methods for visualizing and understanding user preferences, we propose a probabilistic semantic latent factor model, named *SeeP*, that embeds users, items, and topics of reviews in a 2- or 3-dimensional space for visualization and explanation. Different from traditional latent factor models that need to reduce the dimension of high-dimensional latent vectors for visualization, our proposed model directly learns for each user and each item a latent vector in a 2- or 3-dimensional space. Since users and items are in the same visualization space, we propose using distances between users and items, instead of using dot products, to model the ratings. Intuitively, a user close to an item would give a high rating to that item. Moreover, since we want to extract the topics from reviews and use those to explain why a user prefers some items, we also embed review topics into the visualization space. Topics next to a user and an item can give an explanation on why that user gives such a rating to that item. Since our proposed model optimizes and learns the representations of users, items, and review topics at the same time, the learned visualization ensures that both rating behavior and semantic aspects of ratings are preserved as much as possible. This could have important applications for visual analysis of user preferences and exploring item recommendations.



This work is licensed under a Creative Commons Attribution International 4.0 License.

ASONAM '23, November 6-9, 2023, Kusadasi, Turkey

© 2023 Copyright is held by the owner/author(s).

ACM ISBN 979-8-4007-0409-3/23/11.

<http://dx.doi.org/10.1145/3625007.3627486>

In this paper, we make the following contributions:

- We propose a holistic model of user preferences, review topics, and visualization for visualizing and understanding user preferences. Our proposed model, named *SeeP*, embeds users, items, and review topics in a 2- or 3-dimensional space for visualization and explanation. Both user preferences and semantic aspects of ratings are preserved as much as possible in the visualization space. We derive an algorithm to estimate model parameters based on the variational inference approach.
- We conduct extensive experiments on several large real-world datasets. The experimental results show the efficiency and effectiveness of our proposed model.

II. SEMANTIC VISUAL LATENT FACTOR MODEL

A. Problem Definition

The considered problem can be stated as follows: Given N_u users, and N_i items. Let R be the rating matrix where $R_{u,i}$ represents the rating of user u for item i . Let $\mathbf{w}_{u,i}$ be the review by user u for item i . To extract and visualize review topics, since an item i can be described by all reviews given to it, we integrate all reviews for an item i into a single document $\mathbf{w}_i$. Let $\mathcal{D}$ be the set of these item-documents, $\mathcal{D} = \{\mathbf{w}_i\}_{i=1}^{N_i}$, and $\mathcal{V}$ is the vocabulary of the text reviews. We want to find:

- $U = \{\gamma_u\}_{u=1}^{N_u} \in R^{H \times N_u}$ and $I = \{\gamma_i\}_{i=1}^{N_i} \in R^{H \times N_i}$ which are the latent factor matrices for all users and items respectively. Here $H = 2$ or 3 for visualization, and column γ_u and column γ_i represent user u and item i respectively. The distances between users and items reflect the ratings. A user close to an item would give a high rating to that item.
- In addition, our model learns Z latent topics, their visualization coordinates $\Phi = \{\phi_z\}_{z=1}^Z$, and their word distributions $\beta = \{\beta_z\}_{z=1}^Z$. The topic distributions of item-documents are denoted as $\Theta = \{\theta_i\}_{i=1}^{N_i}$. The distances between an item and topics reflect the topic distribution of reviews for that item. Since the distance from a user to an item expresses the rating, we can rely on topics near to that item to explain why that user gives such a rating to the item.

B. Generation and Inference

We model the topic distribution of an item-document i as:

$$\theta_{iz} = p(z|\gamma_i, \Phi) = \frac{\exp\left(-\frac{1}{2}\|\gamma_i - \phi_z\|^2\right)}{\sum_{z'=1}^Z \exp\left(-\frac{1}{2}\|\gamma_i - \phi_{z'}\|^2\right)} \quad (1)$$

here γ_i is the visualization coordinate of item i . It is clear from the equation that we want to keep the topics of item i close to it in the visualization space.

To visualize user preferences toward items, we model the ratings using Euclidean distances between users and items. We assume the following Gaussian distribution over the rating $r_{u,i}$:

$$r_{u,i} \sim \mathcal{N}(f(\gamma_u, \gamma_i), \sigma_r^2) \quad (2)$$

here $f(\gamma_u, \gamma_i)$ is a function of Euclidean distance between user u and user i :

$$f(\gamma_u, \gamma_i) = \mathcal{R} \exp\left(-\alpha^2 \|\gamma_u - \gamma_i\|^2\right) \quad (3)$$

where $\mathcal{R}$ is the max rating that a user could give to an item (i.e., 5 if the 5-point Likert scale is used) and α^2 is a positive scaling factor that will be learned. Intuitively, users give higher ratings to the items that are closer to them in the visualization. Given the above assumptions, the generative process of *SeeP* is as follows:

- 1) For each document $\mathbf{w}_i$ corresponding to item i , $i = 1, \dots, N_i$:
 - a) Draw a document coordinate: $\gamma_i \sim \mathcal{N}(\mathbf{0}, \sigma_i^2 \mathbf{I})$
 - b) For each word $w_{i,m}$ in document $\mathbf{w}_i$:
 - i) Draw a topic: $z \sim \text{Multi}\left(\{p(z|\gamma_i, \Phi)\}_{z=1}^Z\right)$
 - ii) Draw a word: $w_{i,m} \sim \text{Multi}(\beta_z)$
- 2) Given a user u and an item i :
 - a) Draw a rating: $r_{u,i} \sim \mathcal{N}(f(\gamma_u, \gamma_i), \sigma_r^2)$

here γ_i, γ_u are the latent vectors of item i and user u or their visualization coordinates. We treat α , $U = \{\gamma_u\}_{u=1}^{N_u}$, topic visualization coordinates $\Phi = \{\phi_z\}_{z=1}^Z$, and topic word distributions $\beta = \{\beta_z\}_{z=1}^Z$, as model parameters that can be learned by maximizing the following lower bound on the marginal log likelihood:

$$\begin{aligned} \mathcal{L}(\eta, \sigma_i, \sigma_r, U, \Phi, \alpha, \beta; \mathbf{w}_i) &= - \sum_{i=1}^{N_i} D_{\text{KL}}[q(\gamma_i|\mathbf{w}_i, \eta) \| p(\gamma_i|\sigma_i)] \\ &+ \sum_{i=1}^{N_i} E_{q(\gamma_i|\mathbf{w}_i, \eta)} [\log p(\mathbf{w}_i|\gamma_i, \Phi, \beta)] \\ &+ \sum_{u=1}^{N_u} \sum_{i=1}^{N_i} E_{q(\gamma_i|\mathbf{w}_i, \eta)} [\log [\mathcal{N}(r_{u,i}|f(\gamma_u, \gamma_i), \sigma_r^2)]] \end{aligned} \quad (4)$$

here $q(\gamma_i|\mathbf{w}_i, \eta)$ is a variational distribution to approximate the intractable true posterior $p(\gamma_i|\mathbf{w}_i)$. $q(\gamma_i|\mathbf{w}_i, \eta)$ can take a Gaussian form whose parameters are parameterized by an inference neural network:

$$q(\gamma_i|\mathbf{w}_i, \eta) = \mathcal{N}(\gamma_i|f_{\mu_0}(\mathbf{w}_i), f_{\sigma_0^2}^2(\mathbf{w}_i)) \quad (5)$$

here $f_{\mu_0}, f_{\sigma_0^2}$ can be multilayer perceptrons. The expectations in Eq. 4 can be approximated by using a Monte Carlo estimator [12]. More specifically, we sample $\gamma_i^{(l)}$ from $q(\gamma_i|\mathbf{w}_i, \eta)$ by sampling an auxiliary noise sample $\epsilon^{(l)}$ from $\mathcal{N}(0, \mathbf{I})$ and computing $\gamma_i^{(l)} = f_{\mu_0}(\mathbf{w}_i) + \epsilon^{(l)} \cdot f_{\sigma_0^2}(\mathbf{w}_i)$. The expectations can then be approximated as:

$$\begin{aligned} E_{q(\gamma_i|\mathbf{w}_i, \eta)} [\log p(\mathbf{w}_i|\gamma_i, \Phi, \beta)] &\approx \frac{1}{L} \sum_{l=1}^L \log p(\mathbf{w}_i|\gamma_i^{(l)}, \Phi, \beta) \end{aligned} \quad (6)$$

Dataset	#Users	#Items	#Reviews	#Classes
Software	1818	797	11732	15
Arts Crafts and Sewing	55964	22919	431443	10
Music Instruments	27449	10615	216105	14
Yelp Tucson	113974	9250	379719	743
Office Products	101065	27947	727635	3
Industrial and Scientific	10984	5321	70612	25

TABLE I: Dataset statistics

$$\begin{aligned}
E_{q(\gamma_i|\mathbf{w}_n, \eta)} [\log [\mathcal{N}(r_{u,i}|f(\gamma_u, \gamma_i), \sigma_r^2)]] \\
\approx \frac{1}{L} \sum_{l=1}^L \log [\mathcal{N}(r_{u,i}|f(\gamma_u, \gamma_i^{(l)}), \sigma_r^2)] \quad (7) \\
\approx -\frac{1}{2\sigma_r^2} \sum_{u=1}^{N_u} \sum_{i=1}^{N_i} (r_{u,i} - \alpha d(\gamma_u, \gamma_i))^2
\end{aligned}$$

Plug these approximated expectations into Eq. 4, we have a differentiable objective function which can be optimized by stochastic gradient descent¹. The main steps of the inference algorithm are described in Algorithm 1.

III. EXPERIMENTS

A. Datasets and Comparative Methods

We use several real-world datasets from different on-line product and review platforms: Amazon², Yelp³. We adopt 5 product categories from Amazon with different sizes: Software, Arts Crafts and Sewing, Music Instruments, Office Products, and Industrial and Scientific. For Yelp dataset, we keep only reviews for all businesses in Tucson city. The preprocessing step including removal of all duplicate reviews, and following [13]–[15], we use the 5-core version of these datasets in which every user and item has less than 5 reviews are removed. The statistics of all datasets are shown in Table I. The last column shows the number of classes for items in each dataset. We randomly use 90% of each dataset for training and the rest for testing. We only use the text reviews in the training set for training. For text preprocessing, we remove stopwords, stem words, and keep the 5000 most frequent words as the vocabulary.

We compare *SeeP* with models that do not utilize textual reviews including Probabilistic Matrix Factorization⁴ (*PMF*) [16], Non-negative Matrix Factorization⁴ (*NMF*) [17], Neural Collaborative Filtering⁵ (*NCF*) [18], *DGCF*⁶ [19]. For models that leverage textual reviews, we compare *SeeP* with *HFT*⁷ [3], *TransNet-Ext*⁸ [5], *NARRE*⁹ [20], *RGCL*¹⁰ [21], *PLSV*¹¹ [10],

and *LDA*¹² [22]. For *PLSV*, γ_i is learned by using only item-documents. For *LDA*, we extract topic proportions of item-documents and γ_i is then obtained by embedding these topic proportions into the visualization space using *t*-SNE [23]. Finally, with γ_i fixed, we learn γ_u by *NCF* trained on ratings. For visualization, γ_u and γ_i are learned with two dimensions for all methods in our experiments. The hyperparameters of methods are set based on the original implementations. All experimental results are averaged across 5 separate runs.

B. Visualization and Rating Prediction

In this section, we evaluate the performance of our method based on both rating prediction and visualization quality. For rating prediction, Mean Squared Error (MSE) is a widely used metric. γ_u and γ_i in a good visualization should produce a low MSE and thus we can visually explore the items that a user may like. For visualization quality, we also rely on labels of items for evaluation. The intuition is that a good visualization should group items of the same label together. Therefore, we calculate *k*-nearest neighbors (*k*-NN) accuracy in the visualization space. A higher *k*-NN accuracy means a better visualization.

We plot *k*-NN vs. MSE in Figures 1 to demonstrate how methods balance between *k*-NN and MSE for visualization and rating prediction tasks. Figure 1 reports results with *k* = 5 and *Z* = 10. As we can see in the figures, *TransNet-Ext*, *NARRE*, and *RGCL* often have the lowest MSEs because they utilize textual reviews for improving the recommendation. There is a trade-off between *k*-NN accuracy and MSE. Recommendation methods such as *NCF*, *DGCF*, *TransNet-Ext*, *NARRE*, *RGCL*, *HFT*, *PMF*, and *NMF* have low MSEs (i.e., they perform well in the rating prediction task). However, they do not produce good visualizations, indicated by very low *k*-NN accuracies, because these models are not for visualization. In contrast, *PLSV* and *LDA* have high *k*-NN accuracies because they visualize the items based on topic models trained on textual reviews, which can group well the items of the same label. However, since they are not recommendation models, they do not perform well in the rating prediction task, indicated by their very high MSEs. For our proposed model, *SeeP*, it balances well the *k*-NN accuracy and MSE where it has low MSEs and high *k*-NN accuracies in most settings. By modeling ratings, reviews, and visualization in a unified model, *SeeP* can achieve a good performance in most of the datasets in terms of visualization and rating prediction. This shows that the visualization by *SeeP* preserves well user preferences, which is useful for visually exploring and making sense of user rating behavior. Visualization examples in Figures 3 will demonstrate this point further.

C. Topic Coherence

In this section, we evaluate the quality of topics generated by different methods. The goal is to show that besides achieving good performance in rating prediction by utilizing topics

¹ We use *L* = 1 in our experiments

² <https://nijianmo.github.io/amazon/>

³ <https://www.yelp.com/dataset>

⁴ <https://surprise.readthedocs.io/en/stable/index.html>

⁵ https://github.com/hexiangnan/neural_collaborative_filtering

⁶ <https://tinyurl.com/github-dgcf>

⁷ <https://cseweb.ucsd.edu/~jmcauley/>

⁸ <https://github.com/winterant/TransNets>

⁹ <https://github.com/chenchongthu/NARRE>

¹⁰ <https://github.com/JarenceSJ/ReviewGraph>

¹¹ https://github.com/dangpnh2/plsv_vae

¹² https://github.com/akashgit/autoencoding_vi_for_topic_models

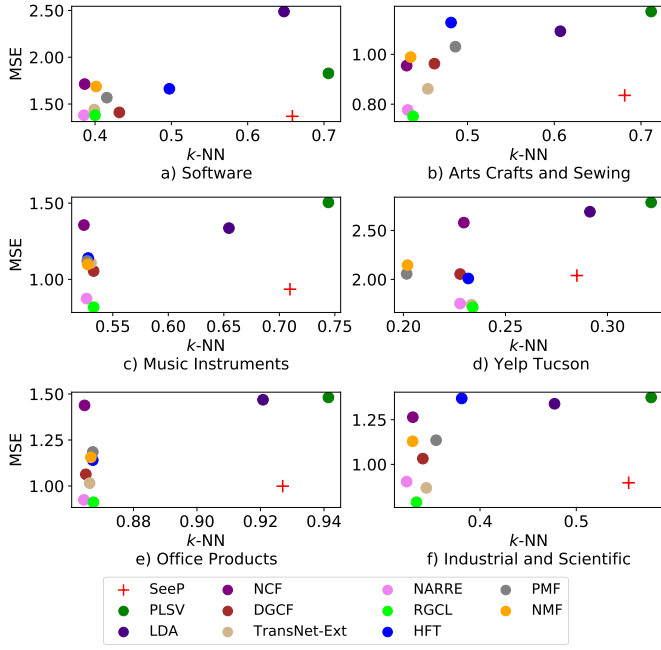


Fig. 1: k -NN vs. MSE across all datasets with $k = 5$, $Z = 10$

in reviews, *SeeP* also gains a competitive performance on topic coherence. Since we use topics to explain the ratings in the visualization, their coherence is very important. To evaluate topic coherence, we use Normalized Pointwise Mutual Information (NPMI) that can be estimated using an external large corpus. For a topic, the NPMI score will be an average over all possible pairs of words of that topic. Figure 2 shows NPMI scores of all methods across different numbers of topics. As shown in the figure, NPMI scores of *SeeP* are comparable to topic model-based methods, *PLSV* and *LDA*. Since, for visualization, *HFT* can only run with two topics, we do not show it in Figure 2.

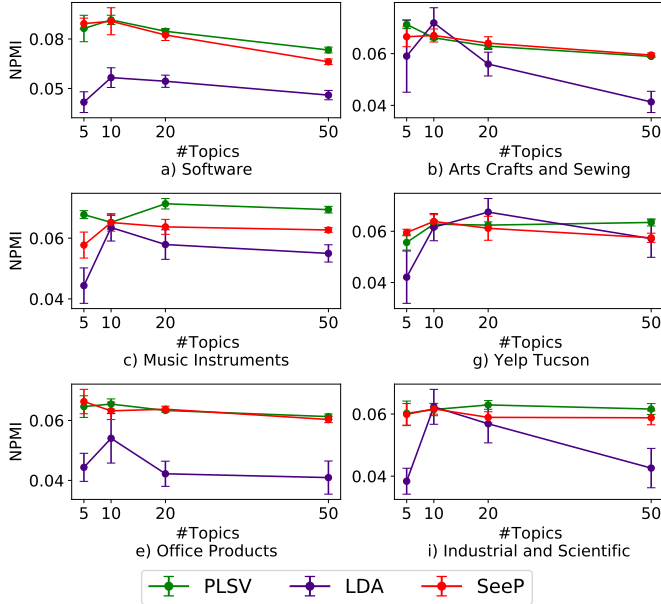


Fig. 2: Topic coherence

D. Visualization Examples

Due to limited space, we show some visualization examples by *SeeP*, *HFT*, and *PLSV* on Arts Crafts and Sewing dataset in Figure 3. In this visualization, each user is represented by a black cross, the colored pluses are items whose colors express the classes of items, and the large red points are topics. In the visualization of *HFT*, the items are mixed up and groups of items are not well-preserved. In contrast, *SeeP* and *PLSV* generate different clusters of items that are consistent with the class labels of items. However, in *PLSV*, all γ_u are grouped close to the center of the visualization, which does not result in a compelling visualization of user-item relationships.

IV. RELATED WORK

In recommender systems, latent factor models are widely used for rating prediction [1], [16]–[18]. These models consider only the numeric ratings and ignore all other sources of information given either by users or items. Recent approaches including deep learning exploit textual reviews to enhance the performance of recommendation systems [5], [6], [20], [21], [24], [25]. One of methods closely related our method is *HFT* where it adds a corpus likelihood based on *LDA* [22] topic models to the traditional latent factor model [3]. *HFT* uses topics as regularizers for more accurately fitting user and item latent factors. None of the above methods can generate a visualization for effectively presenting user preferences and exploring item recommendations.

For visualization in recommender systems, it has shown several important applications in different tasks such as presenting and manipulating user preferences [7], exploring social recommendations [8], and visualizing explanations in recommender systems [9]. In these methods, one possible way for visualization is to embed user and item latent vectors into a 2- or 3-dimensional space using a dimension reduction method such as MDS [7] or t -SNE [23]. However, it is not trivial to incorporate reviews into these methods when performing the dimension reduction of user and item vectors. Some recent document visualization methods can visualize topics in reviews [10], [26]. However, they do not model users, items, and ratings. Therefore, the output visualization may not reflect well user rating patterns. There have been other works for visualization recommender systems which is different from the considered problem in this paper. Visualization recommender systems aim to recommend visualisation methods to users given specific contexts [27], [28].

V. CONCLUSION

We propose a probabilistic semantic latent factor model, named *SeeP*, that embeds users, items, and topics of reviews for visualization and explanation. The user-item interactions are modeled by the distances between users and items in the visualization space. We conduct experiments using several datasets to show the effectiveness of our proposed method in rating prediction, item classification, topic coherence, and generating visualization.

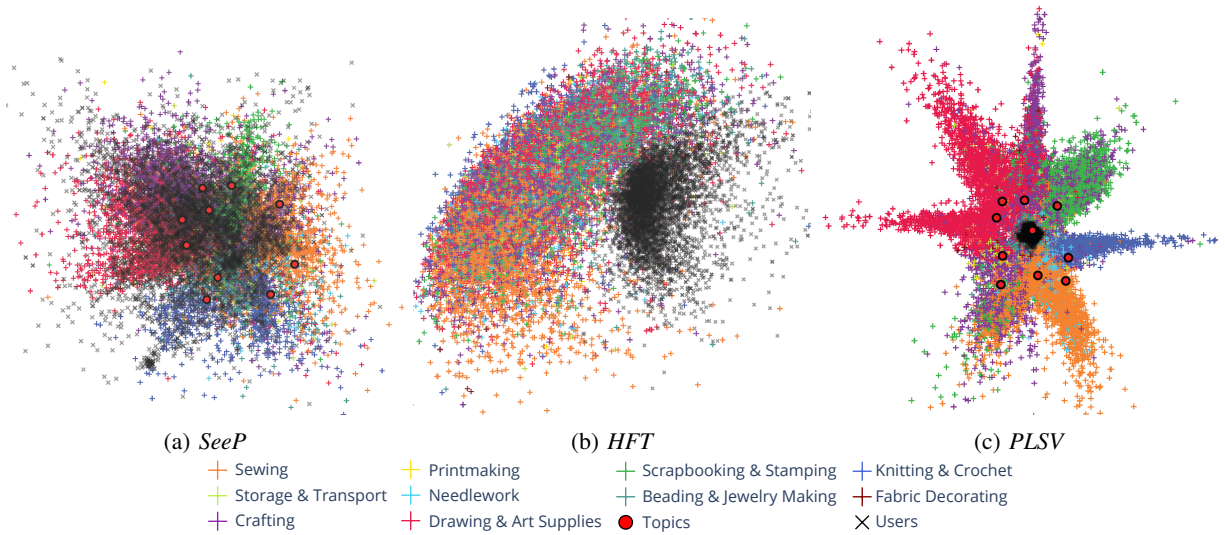


Fig. 3: Visualization of Arts Crafts and Sewing by a) *SeeP*; b) *HFT*; c) *PLSV* ($Z = 10$)

ACKNOWLEDGMENT

This research is sponsored by NSF #1757207 and NSF #1914635.

REFERENCES

- [1] Y. Koren, R. Bell, and C. Volinsky, "Matrix factorization techniques for recommender systems," *Computer*, vol. 42, no. 8, pp. 30–37, 2009.
- [2] D. Lee and H. S. Seung, "Algorithms for non-negative matrix factorization," *Advances in neural information processing systems*, 2000.
- [3] J. McAuley and J. Leskovec, "Hidden factors and hidden topics: understanding rating dimensions with review text," in *Proceedings of the 7th ACM conference on Recommender systems*, 2013, pp. 165–172.
- [4] X. He, T. Chen, M.-Y. Kan, and X. Chen, "Trirank: Review-aware explainable recommendation by modeling aspects," in *Proceedings of the 24th ACM international conference on information and knowledge management*, 2015, pp. 1661–1670.
- [5] R. Catherine and W. Cohen, "Transnets: Learning to transform for recommendation," in *Proceedings of the eleventh ACM conference on recommender systems*, 2017, pp. 288–296.
- [6] X. Guan, Z. Cheng, X. He, Y. Zhang, Z. Zhu, Q. Peng, and T.-S. Chua, "Attentive aspect modeling for review-aware recommendation," *ACM Transactions on Information Systems (TOIS)*, vol. 37, no. 3, 2019.
- [7] J. Kunkel, B. Loepp, and J. Ziegler, "A 3d item space visualization for presenting and manipulating user preferences in collaborative filtering," in *Proceedings of the 22nd international conference on intelligent user interfaces*, 2017, pp. 3–15.
- [8] C.-H. Tsai and P. Brusilovsky, "Exploring social recommendations with visual diversity-promoting interfaces," *ACM Transactions on Interactive Intelligent Systems (TiiS)*, vol. 10, no. 1, pp. 1–34, 2019.
- [9] M. Z. Al-Taie and S. Kadry, "Visualization of explanations in recommender systems," *Journal of Advanced Management Science Vol*, vol. 2, no. 2, pp. 140–144, 2014.
- [10] T. Iwata, T. Yamada, and N. Ueda, "Probabilistic latent semantic visualization: topic model for visualizing documents," in *Proceedings of the 14th ACM SIGKDD international conference on Knowledge discovery and data mining*, 2008, pp. 363–371.
- [11] A. Chaney and D. Blei, "Visualizing topic models," in *Proceedings of the International AAAI Conference on Web and Social Media*, vol. 6, no. 1, 2012, pp. 419–422.
- [12] D. P. Kingma and M. Welling, "Auto-encoding variational bayes," *arXiv preprint arXiv:1312.6114*, 2013.
- [13] J. Tan, S. Xu, Y. Ge, Y. Li, X. Chen, and Y. Zhang, "Counterfactual explainable recommendation," in *Proceedings of the 30th ACM International Conference on Information & Knowledge Management*, 2021.
- [14] N. Wang, H. Wang, Y. Jia, and Y. Yin, "Explainable recommendation via multi-task learning in opinionated text data," in *The 41st International ACM SIGIR Conference on Research & Development in Information Retrieval*, 2018, pp. 165–174.
- [15] Y. Xian, Z. Fu, H. Zhao, Y. Ge, X. Chen, Q. Huang, S. Geng, Z. Qin, G. De Melo, S. Muthukrishnan *et al.*, "Cafe: Coarse-to-fine neural symbolic reasoning for explainable recommendation," in *Proceedings of the 29th ACM International Conference on Information & Knowledge Management*, 2020, pp. 1645–1654.
- [16] A. Mnih and R. R. Salakhutdinov, "Probabilistic matrix factorization," *Advances in neural information processing systems*, vol. 20, 2007.
- [17] X. Luo, M. Zhou, Y. Xia, and Q. Zhu, "An efficient non-negative matrix-factorization-based approach to collaborative filtering for recommender systems," *IEEE Transactions on Industrial Informatics*, vol. 10, no. 2, pp. 1273–1284, 2014.
- [18] X. He, L. Liao, H. Zhang, L. Nie, X. Hu, and T.-S. Chua, "Neural collaborative filtering," in *Proceedings of the 26th international conference on world wide web*, 2017, pp. 173–182.
- [19] X. Wang, H. Jin, A. Zhang, X. He, T. Xu, and T.-S. Chua, "Disentangled graph collaborative filtering," in *Proceedings of the 43rd international ACM SIGIR conference on research and development in information retrieval*, 2020, pp. 1001–1010.
- [20] C. Chen, M. Zhang, Y. Liu, and S. Ma, "Neural attentional rating regression with review-level explanations," in *Proceedings of the 2018 World Wide Web Conference*, 2018, pp. 1583–1592.
- [21] J. Shuai, K. Zhang, L. Wu, P. Sun, R. Hong, M. Wang, and Y. Li, "A review-aware graph contrastive learning framework for recommendation," in *Proceedings of the 45th International ACM SIGIR Conference on Research and Development in Information Retrieval*, 2022.
- [22] D. M. Blei, A. Y. Ng, and M. I. Jordan, "Latent dirichlet allocation," *Journal of machine Learning research*, vol. 3, no. Jan, 2003.
- [23] L. Van der Maaten and G. Hinton, "Visualizing data using t-sne," *Journal of machine learning research*, vol. 9, no. 11, 2008.
- [24] L. Zheng, V. Noroozi, and P. S. Yu, "Joint deep modeling of users and items using reviews for recommendation," in *Proceedings of the tenth ACM international conference on web search and data mining*, 2017, pp. 425–434.
- [25] Y. Tay, A. T. Luu, and S. C. Hui, "Multi-pointer co-attention networks for recommendation," in *Proceedings of the 24th ACM SIGKDD international conference on knowledge discovery & data mining*, 2018.
- [26] D. Pham and T. Le, "Auto-encoding variational bayes for inferring topics and visualization," in *Proceedings of the 28th International Conference on Computational Linguistics*, 2020, pp. 5223–5234.
- [27] Y. Zeng, X. Deng, X. Chen, and R. Jin, "A prediction-oriented optimal design for visualisation recommender systems," *Statistical Theory and Related Fields*, vol. 5, no. 2, pp. 134–148, 2021.
- [28] X. Qian, R. A. Rossi, F. Du, S. Kim, E. Koh, S. Malik, T. Y. Lee, and J. Chan, "Learning to recommend visualizations from data," in *Proceedings of the 27th ACM SIGKDD Conference on Knowledge Discovery & Data Mining*, 2021, pp. 1359–1369.