

Personalized Product Ranking in E-commerce with User-Item Network Models

Nianying Li

*Business School, The University of Sydney, Sydney, Australia
nili0398@uni.sydney.edu.au*

Abstract. Personalisation of product rankings in e-commerce is needed because different users have different interests, demands and browsing conditions. A user-item network model can be employed to represent clicks, favourites, additions to shopping carts, ratings and purchases for personalised Top-k ranking in this paper. This review connects PageRank and Personalised PageRank with graph collaborative filtering, multi-behaviour graph learning, temporal self-supervision and industrial pre-ranking systems. According to the above research, different types of behaviour, the order of interaction and time information, graph construction methods, negative sampling, etc., can affect the ranking results. However, there are still many problems such as scarce and noisy implicit feedback, cold start, popularity bias, preference drift, scalability, privacy, fairness and lack of interpretability. In the future, research will continue to be carried out in the field of combining dynamic heterogeneous graphs, multi-behavioural objectives, causal and self-supervised learning, privacy-preserving computation and graph foundation models. In short, the user-item network model applies linear algebra, probability theory, graph propagation and representation learning to solve the problem of personalised e-commerce ranking. It is convenient to conduct a comparison of the models and decide which one to use.

Keywords: personalized product ranking, user-item networks, graph recommendation, multi-behavior recommendation, e-commerce

1. Introduction

E-commerce platforms must organise the numerous products reasonably for the users. Since most users can only see a small number of products on the search-result page, recommendation panel, or home feed, the quality of the ranking directly impacts product discovery, conversion and user satisfaction. Global rankings based on sales volume, average ratings and recent popularity can provide some reference points; however, they do not consider users' various reasons for choosing a product. Therefore, personalized product ranking can determine the relevance of the candidate products to a particular user and return an ordered list of these products according to their needs. Recently, among the many problems studied in recommender system research, robustness, bias, fairness and evaluation have also been investigated along with ranking accuracy [1].

A network view of users and products can be used to show clicks, favorites, additions to shopping carts, reviews, purchases, etc. These edges can be directed or weighted according to interaction

frequency, recency and behaviour type. This representation allows the ranking model to distinguish weaker signals, such as views or clicks, from the stronger signal of purchase. Recently, some studies have shown that modified PageRank-style propagation can incorporate signed feedback, while personalised restart vectors can be efficiently updated after changes in the network topology [2, 3].

Modern graph-based recommendation is a development based on the old probability propagation method; that is, instead of having fixed transition rules, it learns user and item representations from the structure of a graph. XSimGCL shows that with the addition of a small amount of noise-based contrastive augmentation, the lightweight graph collaborative filtering backbone can be enhanced to improve the accuracy and efficiency of recommendations for sparse user-item graphs and reduce popularity concentration [4]. Therefore, interaction graphs in e-commerce are generally sparse and large, and they consist of many ID and behaviour signals rather than rich node features. At the same time, modern platforms do not observe only one behaviour. A user can view a product many times, save it, add it to the cart, delete it, and purchase it later. A recent journal survey organises multi-behaviour recommendation around data modelling, encoding and training, and suggests that auxiliary behaviours can provide support for a target behaviour such as purchase [5].

Recently, some studies have concentrated on arranging several behaviours according to their meaning and time. A Partial Order Graph Convolutional Network (POGCN) is employed to integrate the correlations and partial orders of various kinds of behaviour into one graph, and LS4SRec builds a global item-transition graph and uses self-supervised disentanglement to model long-term and short-term interests [6, 7]. In light of the above studies, behaviour semantics and temporal context should be regarded as structural signals rather than collapsed into a single binary interaction matrix.

This review focuses on personalized product ranking in e-commerce user-item networks. It does not put forward or realise a new recommender system. Instead, the current study synthesises existing network models, ranking methods, representative studies and practical problems. Section 2 introduces the definitions of e-commerce interaction data, bipartite graphs, matrices, preference vectors and Top-k tasks. Section 3 is about PageRank-style propagation, random walk with restart, graph diffusion, matrix factorization, graph collaborative filtering and ranking evaluation. Section 4 introduces some recent research on data sources, behaviour types, time series, graph design, ranking objectives and evaluation. Section 5 discusses limitations and future directions, which include dynamic graphs, heterogeneous behaviour fusion, cold start, privacy, fairness, explainability and graph foundation models. In short, the purpose is to present the application of linear algebra, probability theory and network learning in individualised commercial ranking.

2. Basic concepts of e-commerce interaction data and user-item networks

The types of e-commerce interaction data are explicit and implicit feedback. Explicit feedback is deliberately given by a user, such as a star rating, written review, like or dislike. It generally has clearer meanings but is relatively scarce. Implicit feedback is the result of normal platform usage, such as impressions, clicks, dwell time, searches, favourites, cart additions, purchases and repeat purchases. It is much more numerous, but the absence of an action does not necessarily mean dislike. A non-clicked product may have been too low on the page, and a clicked product might have aroused curiosity rather than the desire to buy. Therefore, the strength of interaction should not be regarded as a direct indication of interest.

A basic user-item network can be represented as a bipartite graph $G = (U, I, E)$, where U is the set of users, I is the set of products, and E is the set of observed user-item interactions. In the most basic case, there are no user-user and item-item edges. The graph can be represented by an interaction

matrix R where the rows are users and the columns are products. A binary entry indicates whether an interaction occurred, whereas a weighted entry can encode frequency, type of behaviour, etc. For example, the value of a click is lower than that of a purchase, and repeated purchases may be given more weight than a single view. When multiple types of behaviour are modelled separately, the data will be in the form of several matrices or graph layers. An edge can be assigned a behaviour label or a vector of behaviour-specific weights in the case of multi-behaviour graph design.

Some related matrices can be derived from the interaction matrix. The Adjacency Matrix shows which nodes are connected. Degree matrices show the number or total weight of edges connected to each node. Normalising the raw adjacency values produces transition probabilities, so each row or column of the result represents a probability distribution. A preference vector identifies the current user or a small number of seed nodes. Iterative propagation is then used to obtain a ranking vector containing the relevance scores of all products. In the embedding method, the rows and columns of a matrix are turned into low-dimensional vectors, and then the closeness of users and items is determined by means such as an inner product, cosine similarity, a learned scoring function, etc. These operations connect classical matrix factorization with graph propagation because both extract ranking signals from sparse interaction structures.

The output is generally evaluated as a Top-k ranking problem, and the system will place the most relevant products at the beginning of a candidate list. This setting is more suitable for e-commerce product ranking than rating prediction because users are more likely to choose from an ordered list of displayed products than to compare independent numerical estimates of ratings. The task involves data sparsity, uneven strength of feedback, delayed conversion, missing-not-at-random exposure and preference drift. Product popularity also creates a serious structural imbalance, so popular products receive more propagation and training signals. Temporal splitting should therefore be employed to prevent future information from leaking into the training data. A realistic design should specify the observation window, the meaning of a positive edge, the weighting of behaviours, how to handle repeated interactions, and how to construct negative or unobserved examples.

3. Network-based methods for personalized product ranking

Network-based ranking is based on the idea that the relevance of a product is not only due to its own attributes but also because of its position in an interaction structure. Global PageRank is high for nodes that are linked to by many high-ranking nodes. Global PageRank on a user-item network will give too much importance to popular products with a high degree and may not reflect individual interests. Personalised PageRank introduces a user-specific restart vector that keeps the spread of information focused on the target user or a set of seed nodes. At each step, the random walker either takes an outgoing edge or restarts from the target user or seed set. The restart mechanism keeps the ranking local, while multi-step propagation allows indirect paths such as user-product-user-product to affect the score.

Random walk with restart and graph diffusion are related forms. The transition matrix specifies the probability of moving from each node to an adjacent node, and the restart coefficient controls the balance between local selection and the overall network structure. A small restart probability allows information to travel over longer distances, while a large restart probability keeps the recommendation close to the user-specific restart distribution. Behaviour weights can be adjusted so that a purchase contributes more than a click. Time decay can be used to reduce the weight of older interactions, and separate transition matrices can be used to model the transitions from view to cart and from cart to purchase. Therefore, the construction of the graph is a part of the ranking model, not just a neutral preprocessing step.

Matrix Factorisation is used to learn the low-dimensional representations of users and items from the interaction matrix, and based on this idea, Graph Collaborative Filtering can be employed to iteratively aggregate information from neighbouring users and items. Recently, contrastive methods such as XSimGCL have retained a relatively simple graph-propagation structure and used noise-based views to regularise user and item representations, improving long-tail representation and training efficiency [4]. A model of multiple behaviours learns which different behaviours are related to one another. POGCN models partial orders among behaviours and uses a special sampling strategy to show how graph topology and ranking loss can be combined [6].

Recently, both multi-behaviour graph learning and industrial pre-ranking have been studied. Cascading graph contrastive learning (CGCL) forms a cascade of behaviour-specific graphs and improves the representation of user preferences through contrastive learning [8]. InteractRank is a fast two-tower model that has a small cross-interaction term for web-scale pre-ranking [9]. The two studies represent complementary directions rather than opposing approaches.

Ranking Accuracy and Catalog Quality should be judged together. Precision@K, Recall@K, Hit Rate, Mean Average Precision and Normalized Discounted Cumulative Gain (NDCG) can assess the relevance and order of the recommended items; Coverage, Novelty, Diversity and Long-tail exposure show the range and distribution of the recommendation list. Both sets of indicators should be reported so that an increase in accuracy does not conceal an increase in popularity concentration. Offline indicators should also be calculated according to clearly specified negative-sampling and temporal protocols because the same model can appear stronger or weaker when the candidate set changes.

4. State-of-the-art studies of personalized ranking in e-commerce networks

Recent studies can be compared in terms of behaviour data, graph construction, ranking objectives and evaluation settings. The results of different datasets cannot be directly compared due to the fact that rating-based datasets and click-purchase datasets are very different in sparsity, feedback semantics, exposure mechanisms and candidate construction. POGCN models clicks, favourites, cart additions and purchases as partially ordered behaviours, so it forms a joint graph rather than learning independent representations for each behaviour [6]. Its reported deployment on Alibaba demonstrates practical applications of the unified graph, and the method is assessed through public benchmarks and online A/B tests. The study shows that behavioural relations also matter at the industry level, not only in small offline datasets.

LS4SRec separates long-term and short-term interests more clearly [7]. It uses a global item-transition graph, learns independent representations of stable and transient preferences, and reduces noise and sparsity by means of self-supervised objectives. This design addresses the problem that a short browsing period may reflect a specific purpose, a seasonal need or noise. A single static user embedding may be too general for different events, while a pure sequential model may not utilise collaborative signals from other users. LS4SRec combines graph-enhanced sequence representations with explicit long-term and short-term interest modelling.

CGCL is a kind of representation learning that has a cascading graph structure and can capture cross-behaviour dependencies [8]. To improve the representation of users and items, cascading graph propagation and contrastive objectives have been used together. The reported results suggest that explicit connections among different types of behaviour can improve multi-behaviour recommendation. Views, cart additions and purchases are not interchangeable observations in commerce. They often represent successive stages in a funnel, and movement among stages can carry more information than raw frequencies. Although CGCL shows that a cascading graph

structure can improve multi-behaviour recommendation, its effectiveness depends on whether the preset behaviour relationships accurately reflect how users make decisions in different situations and for different kinds of products.

InteractRank addresses a production-oriented ranking problem [9]. Instead of significantly increasing the size of the two-tower model, it adds cross-interaction features from previous interactions and live sequence data. Based on the above experiments, it was found that the online performance was better than both the BM25 baseline and the vanilla two-tower model, and it could be applied at web scale. The main problem with a theoretically expressive model for product ranking is that it may be too slow. Feature summaries of user-query-item interactions can reduce the cost and increase the speed of full-ranking models.

In 2026, the Transition-Aware Graph Attention Network added informative transitions at the level of items, categories and local neighbours to extend multi-behavior sequential recommendation [10]. A sparse structured graph is employed to construct a transition-aware attention mechanism with linear complexity in terms of the length of the sequence. This is a good direction for e-commerce, as the long history will make quadratic attention too costly. The study reports industrial applications and improvements in business indicators, showing that graph sparsification and transition selection can support both accuracy and efficiency.

Recently, other research has also been carried out on the problem of core ranking. MetaGC-MC is a graph convolution and meta-learning method that can solve the cold-start recommendation problem with or without auxiliary information [11]. A review of explainable graph-based recommender systems has been conducted to organise approaches based on learning methods, explanation methods and types, and it has been found that structural paths and subgraphs can provide human-understandable reasons [12]. FedKGRec is a distributed, knowledge-graph-aware recommendation model that applies local differential privacy to local model parameters before they are uploaded to a central server for aggregation [13]. Together, these methods extend personalised ranking beyond accuracy by addressing cold start, privacy and explainability.

Graph Foundation Models are an emerging type of general model. Recently, survey research has introduced graph foundation models as generalizable models that can be applied to many graph problems and areas [14]. For e-commerce, to address the problem of limited interaction data, structural pre-training could be combined with product titles, descriptions, images, brands, categories, reviews and knowledge-graph relations. Recent journal research on multi-behaviour sequential recommendation uses page views, favourites, cart additions, purchases, time intervals and attribute information to learn user preferences from e-commerce interaction sequences [15]. Therefore, advanced foundation-model methods need to be assessed together with existing multi-behaviour graph and sequential baseline methods. Graph Foundation Models also require substantial computation and may not meet the low-latency demands of the ranking pipeline; therefore, their advantages should be compared with those of simple graph and behaviour baselines.

Across the above studies, reported performance depends heavily on experimental design. Changes in the kind of feedback, temporal splitting, negative sampling, size of the candidate set, graph normalisation and popularity distribution can significantly alter the apparent ranking quality of the same model. Therefore, cross-study comparisons should only be carried out after the preprocessing and evaluation protocols have been aligned.

5. Limitations and prospects

The first major limitation is sparsity. A typical user interacts with only a small proportion of the catalog, and purchases are much less frequent than impressions or clicks. Multi-hop paths in graph

propagation can alleviate sparsity, but weak or noisy evidence may also be spread. Adding auxiliary behaviours increases the number of edges, but their meanings are unclear. A click may be accidental; adding an item to the cart may serve as a temporary bookmark; and a purchase may be made for someone else. In the future, models should account for confidence and uncertainty rather than assign the same fixed weight to all behaviour types. Personalized behaviour weighting is necessary because a shopping cart operation may indicate purchase intention for one person but not for another.

Cold start is still a problem; new users and new products have few or no data. MetaGC-MC addresses sparse-interaction cold start using graph-based meta-learning with or without auxiliary information [11]. However, the general problem needs to be better connected with side information. Product text, images, category hierarchies, brand relations, prices, seller information and knowledge graphs are the initial representations before the accumulation of interactions. The context of the session and search queries, as well as any optional preferences of new users, can be used to construct temporary seed vectors. One possible direction is to combine the above signals with Graph Foundation Models and adapt general semantic knowledge to the interaction structure of the platform.

The problem of popularity bias also exists. High-degree products are more likely to receive exposure, attract more clicks and spread further. This may improve conventional accuracy metrics but reduce novelty, catalogue coverage and opportunities for new sellers. Normalisation, inverse propensity weighting, calibrated re-ranking and long-tail-aware losses address the problem in different ways, and each changes the balance between user relevance, seller exposure and platform revenue. In the future, head, torso and tail products should be evaluated separately to determine how ranking quality changes with user activity and product popularity. The evaluation should also distinguish the recommendation quality of already visited products from the discovery of new products.

People's interests differ. A long-term profile will have stable preferences but miss a current shopping need, and a short session will be noisy. LS4SRec and transition-aware methods show how to model the duration and change of interest over time [7, 10]. In the future, dynamic graphs should be updated as nodes, edges and weights change. Time will affect both the decay of the edge and the meaning of the behaviour sequence. For example, adding items to the cart after viewing is different from adding them after several months. A dynamic graph model can also identify seasonal cycles and differentiate temporary desires from fixed ones.

Model scale is a practical limitation. A large platform has millions of users, products and interactions, so ranking must meet strict latency requirements. Full-graph message passing, deep neighbourhood expansion and quadratic sequence attention can be too expensive. Relevant methods include sparse transition graphs, neighbourhood sampling, cached embeddings, incremental updates and multi-stage ranking. InteractRank and a transition-aware graph model show that production systems often enhance interaction modelling without sacrificing retrieval and pre-ranking speed [9, 10]. In the future, memory use, training cost, update frequency and inference latency should also be reported.

Offline assessment has its own problems. Historical logs contain the policy that created them, so unexposed products are regarded as unknown even if a user might have liked them. Metrics that can be calculated from sampled negatives are unstable, and although there is an increase in NDCG, it may not be associated with an improvement in satisfaction, revenue, retention or diversity. Online experiments provide stronger evidence, but they are expensive and difficult to reproduce. Better research practice should include temporal splits, several negative-sampling methods, full-ranking tests where possible, statistical significance tests, calibration, coverage, diversity and subgroup

analysis. Public datasets should retain impression and exposure information to differentiate between preference and display bias.

Personalised ranking increasingly relies on detailed records of people's behaviour. FedKGRec keeps the raw interaction history locally and applies local differential privacy to model parameters before aggregation [13]. Model updates may still contain membership or preference information. In the future, Federated Learning will be used in conjunction with secure aggregation, differential privacy and user-controlled data retention. Privacy needs to be taken into account at the design stage, with explicit analysis of its effect on model accuracy.

Fairness and explainability affect adoption. The ranking system may give some sellers and brands a disproportionately large number of impressions. The interests of consumers and businesses may not align, so the trade-offs among different stakeholders should be made clear. Graph-based explanations can find the influential path, nearby users, past behaviour and product relations, but the quality of the explanation is not determined only by the model structure. According to the explainability review, the learning method, explanation method and explanation type should be evaluated independently [12]. In the future, we hope that the explanation will be consistent with the model, easy for users to understand, and useful for contesting or correcting the recommendation.

A good technical direction should be a combination of dynamic heterogeneous graphs, multi-behaviour learning and rich semantic representation. Heterogeneous graphs can contain users, products, brands, categories, sellers, queries and attributes, etc., and typed edges can be used to retain the meaning of view, cart, purchase, similarity and ownership relations [14]. Self-supervised objectives do not need labels, causal methods can distinguish between preference and exposure, and graph foundation models can provide reusable structural representations for all kinds of tasks and domains. At the same time, simpler PageRank-type methods are still useful because they are more intuitive, relatively easy to expand and update, etc. A rich research agenda should not be based on the assumption that increased model complexity is always better. Future research should determine what structural information to extract, how much temporal detail is needed, and what semantic context produces measurable gains under realistic latency, privacy and fairness constraints.

In general, future research should move away from using a single indicator for recommendation and adopt a system perspective. The rank of personalised products should be based on the following: relevance, newness, diversity, seller exposure, computational cost, robustness, privacy and user control. The user-item network is still a practical organisational abstraction because all of the above issues can be represented in terms of nodes, edges, weights, propagation rules, constraints and evaluation criteria. The problem is how to build a model that preserves mathematical clarity while meeting the behavioural and operational demands of today's e-commerce.

6. Conclusion

This study reviews personalised product ranking in e-commerce based on user-item networks. Interaction events such as views, favourites, cart additions, ratings and purchases can be expressed in the form of sparse matrices, weighted bipartite graphs, transition probabilities, user-specific preference vectors, etc. Classical PageRank and Personalised PageRank are the basic concepts of network propagation, and matrix factorization and graph collaborative filtering are used to learn compact representations of users and products from large interaction structures. Recently, based on the above foundation, some new methods have been put forward, such as multi-behaviour graphs, temporal self-supervision, cascading graph contrastive learning, transition-aware attention and efficient industrial pre-ranking. Together, these methods show that the ranking performance is not

only a result of model architecture but also depends on how behaviour, time, graph edges, candidate sets and negative samples are defined.

Some problems have not been solved in the studies above. Sparse and noisy implicit feedback can weaken learned preferences; cold-start users and products have few graph connections, and high-degree products may be displayed too often. Other unresolved problems include interest fluctuations, large-scale computation, privacy risks, unfairness across users and sellers, and limited explainability. Therefore, in the future, research should integrate dynamic heterogeneous graphs with robust multi-behaviour models, extensive product information, causal and self-supervised objectives, privacy-preserving learning, calibrated re-ranking, etc. These developments should be evaluated using temporally realistic data divisions, transparent candidate construction, various ranking and catalog quality indicators, efficiency analysis and, where possible, online tests.

In short, the user-item network model combines linear algebra, probability, graph propagation and representation learning to solve the problem of personalised ranking. Its practical value depends on how ranking accuracy is balanced with newness, diversity, computational cost, robustness, privacy, fairness and user control. A good ranking system for e-commerce should be judged as a complete decision-making process, not just as a single predictive model.

References

- [1] Li, Y., Liu, K., Satapathy, R., Wang, S., & Cambria, E. (2024). Recent developments in recommender systems: A survey [review article]. *IEEE Comput. Intell. Mag.*, 19(2), 78-95. <https://doi.org/10.1109/MCI.2024.3363984>
- [2] Gu, K., Fan, Y., & Di, Z. (2022). Signed PageRank on online rating systems. *J. Syst. Sci. Complex.*, 35(1), 58-80. <https://doi.org/10.1007/s11424-021-0124-2>
- [3] Bautista, E., & Latapy, M. (2022). A local updating algorithm for personalized PageRank via Chebyshev polynomials. *Soc. Netw. Anal. Min.*, 12(1), Article 31. <https://doi.org/10.1007/s13278-022-00860-5>
- [4] Yu, J., Xia, X., Chen, T., Cui, L., Hung, N. Q. V., & Yin, H. (2024). XSimGCL: Towards extremely simple graph contrastive learning for recommendation. *IEEE Trans. Knowl. Data Eng.*, 36(2), 913-926. <https://doi.org/10.1109/TKDE.2023.3288135>
- [5] Kim, K., Kim, S., Lee, G., Jung, J., & Shin, K. (2026). A comprehensive survey on multi-behavior recommender systems: Extended taxonomy and recent advances. *Int. J. Data Sci. Anal.*, 22(1), Article 40. <https://doi.org/10.1007/s41060-025-01019-z>
- [6] Zhang, Y., Bei, Y., Chen, H., Shen, Q., Yuan, Z., Gong, H., Wang, S., Huang, F., & Huang, X. (2024). Multi-behavior collaborative filtering with partial order graph convolutional networks. In *Proceedings of the 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining* (pp. 6257-6268). ACM. <https://doi.org/10.1145/3637528.3671569>
- [7] He, L., Yan, W., Yi, T., & Sun, H. (2026). Self-supervised graph neural sequential recommendation with disentangling long- and short-term interest. *ACM Trans. Recomm. Syst.*, 4(1), Article 10, 1-25. <https://doi.org/10.1145/3723173>
- [8] Yang, J., Li, X., Li, B., Tian, L., Xu, B., & Chen, Y. (2024). Cascading graph contrastive learning for multi-behavior recommendation. *Neurocomputing*, 610, 128618. <https://doi.org/10.1016/j.neucom.2024.128618>
- [9] Khandagale, S., Juneja, B., Agarwal, P., Subramanian, A., Yang, J., & Wang, Y. (2025). InteractRank: Personalized web-scale search pre-ranking with cross interaction features. In *Companion Proceedings of the ACM Web Conference 2025* (pp. 287-295). ACM. <https://doi.org/10.1145/3701716.3715239>
- [10] Jin, H., Yang, G., Chan, Z., Yuan, Y., Li, L., Sun, F., Yang, Y., Wu, J., Jiang, Y., & Zheng, B. (2026). Multi-behavior sequential modeling with transition-aware graph attention network for e-commerce recommendation. In *Proceedings of the ACM Web Conference 2026* (pp. 8681-8684). ACM. <https://doi.org/10.1145/3774904.3792939>
- [11] Shu, H., Chung, F.-L., & Lin, D. (2023). MetaGC-MC: A graph-based meta-learning approach to cold-start recommendation with/without auxiliary information. *Inf. Sci.*, 623, 791-811. <https://doi.org/10.1016/j.ins.2022.12.030>
- [12] Markchom, T., Liang, H., & Ferryman, J. M. (2026). Review of explainable graph-based recommender systems. *ACM Comput. Surv.*, 58(6), Article 138, 1-35. <https://doi.org/10.1145/3772273>

- [13] Ma, X., Zhang, H., Zeng, J., Duan, Y., & Wen, X. (2024). FedKGRec: Privacy-preserving federated knowledge graph aware recommender system. *Appl. Intell.*, 54(19), 9028-9044. <https://doi.org/10.1007/s10489-024-05634-4>
- [14] Liu, J., Yang, C., Lu, Z., Chen, J., Li, Y., Zhang, M., Bai, T., Fang, Y., Sun, L., Yu, P. S., & Shi, C. (2025). Graph foundation models: Concepts, opportunities and challenges. *IEEE Trans. Pattern Anal. Mach. Intell.*, 47(6), 5023-5044. <https://doi.org/10.1109/TPAMI.2025.3548729>
- [15] Chen, Y., Cao, Q., Huang, X., & Zou, S. (2024). Multi-behavior collaborative contrastive learning for sequential recommendation. *Complex Intell. Syst.*, 10(4), 5033-5048. <https://doi.org/10.1007/s40747-024-01423-1>