

Measuring the Ranking Quality of Recommendations in a Two-Dimensional Carousel Setting

Nicolò Felicioni
Politecnico di Milano, Italy
nicolo.felicioni@polimi.it

Fernando B. Pérez Maurera
ContentWise
fernando.perez@contentwise.com

Maurizio Ferrari Dacrema
Politecnico di Milano, Italy
maurizio.ferrari@polimi.it

Paolo Cremonesi
Politecnico di Milano, Italy
paolo.cremonesi@polimi.it

ABSTRACT

Movie-on-demand and music streaming services usually provide the user with multiple recommendation lists, i.e., carousels, in a two-dimensional user interface, each generated according to different criteria (e.g., TV series, popular artists, etc.). In this two-dimensional setting it is not appropriate to use traditional ranking metrics designed for a single ranking list. It is well known that users do not explore a two-dimensional interface one row at a time, but rather focus their attention in a triangular area at the top-left corner. Furthermore, it is frequent for user interfaces to hide some items or lists due to space constraints, which can be shown by performing certain actions (i.e., click, swipe). In this paper we extend the widely used NDCG to a two-dimensional recommendation setting with a formulation that allows to account both the two-dimensional user exploration behaviour and interface-specific design. We also compare the proposed extension against single-list NDCG highlighting that they can lead to a different choice of the optimal algorithm in offline evaluation.

CCS CONCEPTS

• Information systems → Collaborative filtering; Recommender systems; • General and reference → Evaluation.

KEYWORDS

Recommender Systems; User Interface; Evaluation

ACM Reference Format:

Nicolò Felicioni, Maurizio Ferrari Dacrema, Fernando B. Pérez Maurera, and Paolo Cremonesi. . Measuring the Ranking Quality of Recommendations in a Two-Dimensional Carousel Setting. In *International Workshop on Industrial Recommendation Systems at KDD '21*, August 15, 2021, Online. ACM, New York, NY, USA, 6 pages.

1 INTRODUCTION

Traditionally, in the Information Retrieval and Recommender Systems domains, the objective has been to provide the user with the best possible ranked list of results [6, 11, 22]. For this reason, many metrics were developed to evaluate the quality of a one-dimensional ranked list. A common assumption is that users will navigate the list according to its order, therefore it is better for a correct recommendation to be at the beginning of the list.

There are however several scenarios that do not fit into these assumptions, mainly when the results are presented in a two-dimensional grid rather than a single list. This is true both in information retrieval [3] and in recommendation systems applications, in particular for video-on-demand streaming services [7, 18, 26] and music streaming platforms [1, 9]. Those services usually provide users with multiple rows of thematically coherent recommendations (e.g., the most popular movies, a specific genre, new releases, and so on, see Figure 1). These rows are referred to as *widgets*, *shelves* or as *carousels*.

A simple way to adapt one-dimensional ranking metrics to a two-dimensional interface is to concatenate all recommendation lists into a single one. This strategy does not make realistic assumptions and, we argue, is not appropriate. First, it is known that users do not explore each carousel sequentially from the first to the last, as concatenating them assumes. Rather, users start from the top-left corner of the screen and proceed to explore the items both to the right and to the bottom [16, 27]. This effect is also known as "golden triangle" or "F-pattern". A visual example from an information retrieval application [3] is shown in Figure 2. Another example from a video streaming service [18] is shown in Figure 3. In addition to this user behaviour, many websites and mobile applications present carousels that are *swipeable* [1], i.e., the user can swipe horizontally or vertically to reveal more items as well as lists that were not previously visible. This is a common way to overcome the limited space available in the user interface allowing to fit more

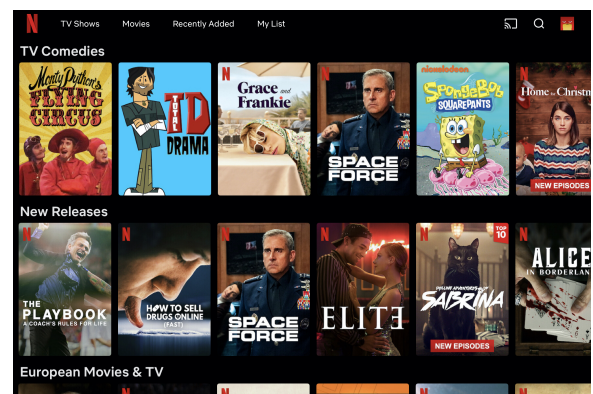


Figure 1: The Netflix homepage, an example of carousel user interface in the multimedia streaming domain.

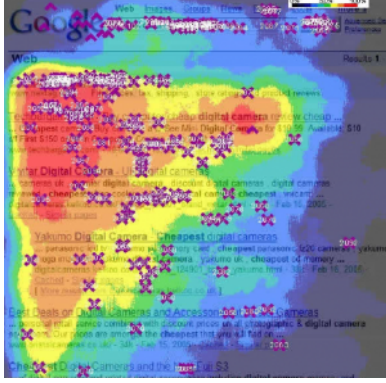


Figure 2: When using a search engine users concentrate their attention on the top-left corner (golden triangle) [3].

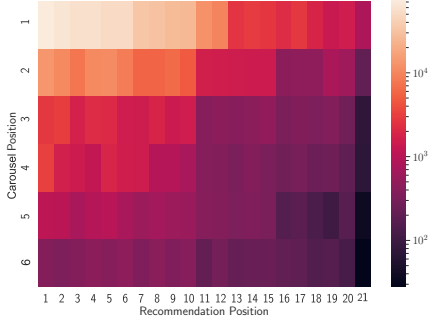


Figure 3: Visualization of the number of user interactions on each position on a user interface [18]. A demarcation between the first and second half of the columns is visible.

recommendations and carousels that the user can easily browse. However this puts additional overhead on the user that has to actively interact with the system to access the recommendations. Hence, it is preferable for a correct recommendation to be visible with the least possible number of user actions, as also noticed in [15].

In order to take those factors into account, in this paper we propose to extend the one-dimensional NDCG metric to consider both the two-dimensional user exploration behaviour and the user interface characteristics. We show that the two metrics can lead to different results when used to select which recommenders to use in the carousel interface.

The rest of the paper is organized as follows, in Section 2 we summarize the characteristics of a carousel setting, in Section 3 we formulate an extended version of NDCG, in Section 4 we perform an offline comparison of the results in a single list and carousel interface. Finally in Section 5 we draw the conclusions.

2 CHARACTERISTICS OF A CAROUSEL SETTING

The carousel interface layout and the way it is usually generated by video-on-demand and music streaming platforms has important characteristics that distinguish it from a single-list setup:

Interface: A two dimensional user interface with multiple carousels. Some carousels or recommendations may be hidden due to limited page size and be accessible only via user actions (i.e., click, swipe).

Recommendations: The lists shown to the users are generated with different algorithms or by different providers and no single post-processing step is applied. While each individual recommendation list does not contain duplicates, the same item may appear in multiple carousels.¹²

User Behaviour: The user will focus on the top-left triangle of the screen rather than exploring the carousels sequentially. Furthermore, they will explore the recommendations in different ways according to which actions they need to perform in order to reveal them.³

While a carousel layout may seem similar to a traditional merge-list embedding, where multiple recommendation list are combined into one, this is not the case. In a real scenario, there are multiple constraints. First, the carousels may be generated by different content providers, each of them unaware of how the other lists are generated or by whom. This means that the composition of the layout as well as the recommendations of the other providers are, in general, not known. It is for this reason that different carousels may contain similar recommendations. Furthermore, a content provider that wishes to select the optimal carousels to display has limited degrees of freedom and can only alter the content and relative ordering of those it is tasked to provide.

3 EXTENDING ONE-DIMENSIONAL NDCG

One of the most used metrics for ranked list evaluation is the *Discounted Cumulated Gain* (DCG), as well as its *Normalized* version (NDCG) [13, 14]. This metric comes from the information retrieval domain and is widely used to evaluate recommendation systems. The DCG metric relies on two assumptions:

- (1) highly relevant results are more valuable for a user;
- (2) within a list of results, it is preferable to have relevant results in the first positions

Let c be the recommendation list length, i.e., cutoff, and $rel(i)$ the relevance of the item in position i . The DCG is defined as the following discounted sum of gains:

$$DCG = \sum_{i=1}^c gain(i) \cdot discount(i)$$

The *gain* function is responsible for rewarding highly relevant results, while the *discount* function introduces a penalization that

¹A significant example are *content aggregators*, which combine carousels from different providers: Netflix, Youtube, Prime Video, etc.

²For example, in the Netflix homepage shown in Figure 1 the TV series *Space Force* appears both in the *TV Comedies* and *New Releases* carousels.

³Usually users tend to navigate more easily with simple swipes rather than repeated mouse clicks, hence their behaviour, as it is known, will change according to the device.

should increase the further the item is from the beginning of the list.

One of the most used formulations for the DCG is the following [2]:

$$DCG = \sum_{i=1}^c \frac{2^{rel(i)} - 1}{\log_2(i+1)}$$

Hence, $gain(i) = 2^{rel(i)} - 1$ and $discount(i) = \frac{1}{\log_2(i+1)}$. Notice that this formulation is only one of many possible formulations for the DCG. Several other ways of rewarding and discounting results have been proposed in previous research [17, 28]. In the following, we will start from this formulation and extend it since it is one of the most used. Other types of gain and discount functions can be extended in an analogous way. We leave the analysis of different gains and discounts as future work.

In a two-dimensional scenario, the standard DCG definition could be naively adapted in the following way. Let h be the horizontal dimension of the interface (i.e., the length of each carousel) and v the vertical dimension of the interface (i.e., the number of carousels). The carousels can be concatenated in a single list of length $c = v \cdot h$ items on which the standard DCG formulation can be applied. This strategy assumes that the users will explore all carousels sequentially, from the first to the last, which, as previously discussed, is not consistent to the user behaviour and does not account for the interface navigation constraints. Therefore, we suggest researchers *do not* apply this strategy as it does not represent a realistic scenario.

Thus, inspired by [13], we make the following assumptions the two-dimensional DCG should meet:

- (1) highly relevant results are more valuable for a user;
- (2) a relevant result is valuable to the user only when it is first seen;
- (3) within a grid of results, it is preferable to have relevant results close to the top-left corner
- (4) it is preferable that relevant items are immediately visible to the user or can be made visible with few user actions

In order to account for this set of assumptions, we propose to extend the metric in the following way:

$$2DCG = \sum_{i=1}^v \sum_{j=1}^h gain(i, j) \cdot discount(i, j)$$

As in the one-dimensional version, the *gain* function is responsible for rewarding highly relevant results, according to assumptions (1) and (2). The *discount* function, instead, should account for the penalty related to the position and number of user actions, according to assumptions (3) and (4).

Inspired by the one-dimensional version, we fix $gain(i, j) = 2^{rel(i, j)} - 1$. Instead, the *discount* will depend on the position in the layout, allowing ample freedom on how to define it in different use cases.

The normalized version of this metric, $N2DCG$ will be defined as $N2DCG = 2DCG/I2DCG$. $I2DCG$ will be the 2DCG of the *ideal ranking*. In a single list setting the ideal ranking is the list which contains the relevant items in decreasing relevance from the beginning of the list. In the generalized two-dimensional layout it contains the user's most relevant items, ranked according

to decreasing relevance in positions with decreasing position discount. The ideal ranking meets the following constraints: for any pair of cells $(i, j), (k, l)$ of the matrix, $gain(i, j) \geq gain(k, l)$ if $discount(i, j) > discount(k, l)$.

Relevance. As stated in assumption (2), a relevant item is valuable for the user only when it is first encountered. This means that if a relevant item appears multiple times, each in a different carousels, it should be considered as relevant only in its *best* position. We define such position as the one with the highest *discount*. Function $rel(i, j)$ should be modified accordingly.

Single List Discount. It is possible to represent in this formulation the traditional single list DCG by calculating the position of cell in coordinates i, j if all carousels lists would be concatenated:

$$discount_{singleList}(i, j) = (\log_2((i-1) \cdot h + j + 1))^{-1}$$

As previously mentioned, this formulation is not grounded in a realistic scenario because it does not reflect the user behaviour (see Figure 4a), therefore we argue it should not be applied.

Golden Triangle Discount. In order to account for the *golden triangle* behaviour, as per assumption (3), the position discount should decrease as the distance of the cell from the top-right corner increases:

$$discount_{triangle}(i, j) = (\log_2(\alpha \cdot i + \beta \cdot j))^{-1}$$

The coefficients α, β are two weights that can be used to account for different types of user behaviors. For instance, let us assume a scenario where users are more inclined to explore the vertical dimension. In this case, α should be set to a low value in order to penalize less the vertical dimension. In order to make the discount start from 1, α and β should be ≥ 1 since the base of the logarithm used is 2. Notice that this is true only because we are extending a logarithmic discount function. For other discount functions [17, 28] the constraints can change.

The resulting discount is shown in Figure 4b (we set $\alpha = \beta = 1$ for simplicity).

User Actions Discount. Lastly, in order to account for assumption (4) the position discount should decrease the more actions are required by the user to make that position visible. In a carousel interface there is an initial rectangular portion of the recommendations that are immediately shown to the user. We refer to the number of items visible as $init_h$ and to the number of carousels visible as $init_v$, see Figure 5. In order to reveal more items, the user needs to perform a certain action, i.e., click on a desktop, swipe on mobile devices. Each of these actions will reveal a certain number of new items within the currently visualized recommendation lists. Different platforms and devices will correspond to different *swipe steps*, i.e., the number of items that will be revealed after a single swipe. We will call this quantity $step_h \in \{1, 2, \dots, init_h\}$. For example, on Netflix every click will replace all items displayed on the clicked carousel, in which case $step_h = init_h$. The same principle holds for the vertical dimension, where the user can navigate performing actions that will each display $step_v$ new carousels.

Based on this definition, we now add to the triangle penalty a term to account for the number of actions that the user will need to perform in order to visualize the item. To do so we define some

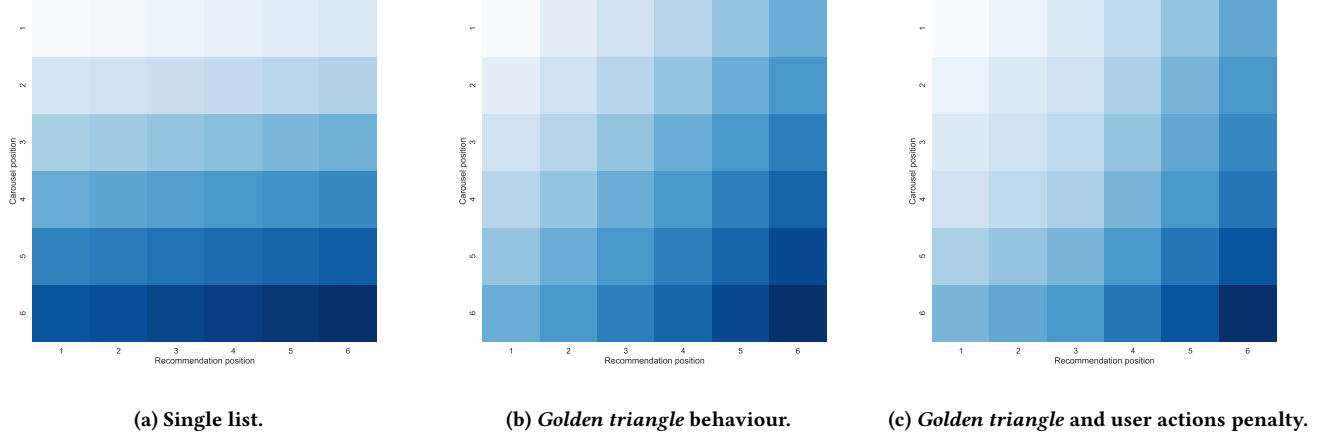


Figure 4: A visual comparison of the two-dimensional penalty function under different assumptions. Figure 4a refers to carousels concatenated in a single list. The other figures refer to the two-dimensional penalty which accounts for the *golden triangle* behaviour only, see Figure 4b as well as the number of user actions, see Figure 4c.

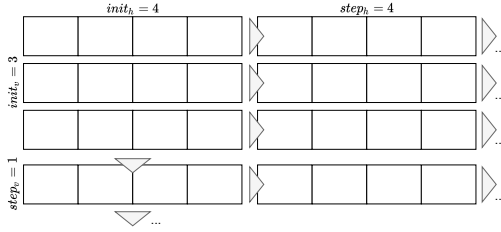


Figure 5: An example interface where 3 carousels, with 4 items each, are visible. A horizontal swipe reveals 4 items, while a vertical swipe reveals one additional carousel.

auxiliary functions. The first one is used to check whether at least a user action, i.e., swipe, is needed to visualize that item in a certain position p given that the interface initially shows $init$ positions:

$$\text{isSwipeNeeded}(p, init) = \begin{cases} 1, & \text{if } p - init > 0 \\ 0, & \text{otherwise} \end{cases}$$

Then, we define a function to count the number of actions needed to visualize an item, given that each action shows $step$ positions:

$$\text{swipes}(p, init, step) = \text{isSwipeNeeded}(p, init) \cdot \left\lceil \frac{p - init}{step} \right\rceil$$

In the particular case where $init = step$, calculating the number of swipes becomes simpler:

$$\text{swipes}(p, step) = \left\lceil \frac{p}{step} \right\rceil$$

The final discount will account for both the triangle discount and the number of user actions, as previously defined:

$$\text{discount}_{\text{actions}}(i, j) = (\log_2(\alpha \cdot i + \beta \cdot j + \gamma \cdot \text{swipes}(i, init_v, step_v)) + \lambda \cdot \text{swipes}(j, init_h, step_h))^{-1}$$

Notice that this formulation accounts for both vertical and horizontal swipes. The coefficients $\alpha, \beta, \gamma, \lambda$ are four positive weights that can be used to account for different types of user behaviors. The first two weights (α and β) control the general penalization of the vertical and horizontal dimensions, respectively. As we previously said, they should be ≥ 1 in order for the total discount to start from 1. Controlling γ and λ , instead, it is possible to penalize more or less the user actions needed to reveal a certain item. For example, it could be that items presented together in the same carousel have a similar probability of interaction (see the first 10 elements of the first carousel in Figure 3). Hence, the horizontal dimension should be penalized less. Another possibility is that, on a desktop device, the horizontal swipe done with a mouse click will have a higher weight than the same swipe done with a touch on a mobile device.

For illustrative purposes, let us consider a possible scenario for a mobile device, where the screen contains 4 carousels and 3 recommendations each. We set the horizontal and vertical steps to 1, $\alpha, \beta, \gamma, \lambda$ are set to 1 as well. The resulting discount is shown in Figure 4c.

4 EXPERIMENTS

In this section we provide an example of the different behaviour of NDCG and N2DCG in an offline experimental scenario. We consider a setting where given a set of recommendation models and a certain number of carousels, the goal is to select which models to use to generate each carousel. We show that the two metrics yield to different carousel layouts. In order to represent a scenario where a carousel interface would be used, we selected the widely known movie recommendations dataset *MovieLens10M* dataset [10], containing 70k users, 10k items and 10M ratings.

The set of models that can be selected, i.e., M , contains several simple and widely known models that have shown to provide competitive results in recent evaluations [8]. For *Non-Personalized* models we selected a TopPopular recommender. As *KNN* models

we included ItemKNN [23] and UserKNN [24], both computing the similarity with cosine and shrinkage. We included the *Graph-based* models $P^3\alpha$ [5] and $RP^3\beta$ [20], which define a bipartite graph of users and items and simulate a random walk. We added various *Matrix Factorization* models, some developed for explicit interactions: PureSVD [6], FunkSVD [8] and Non-negative MF (NMF) [4]; as well as others developed for implicit interactions: MF BPR [21], IALS [12]. We included the widely known *Item-Based machine learning* models SLIM [19], SLIM BPR and the more recent $EASE^R$ [25]. Finally, we included the *Content-based* model ItemKNN CBF, which computes the item similarities from item features. using cosine similarity with shrinkage.

We split the data by randomly selecting 80% of interactions for the training set and 10% for validation and test set. Each model was optimized on the validation data, following the best practices and value ranges reported in [8], using a Bayesian search with 50 cases.

Since the purpose of this paper is not to propose an algorithm for the selection of carousels but to show that the two metrics lead to different results, we rely on a simple greedy strategy. At the beginning the page is empty and all candidate algorithms are evaluated independently on the validation data. The model with the best recommendation quality is selected as first carousel. The process repeats for the following carousels, however, in this case, the candidate model will be evaluated by taking into account all the *previous* carousels. According to the definition of relevance provided in Section 3, a correct recommendation of an item by the candidate model may overtake another of the same item in a previous carousels if it has a better position discount. For example, a correct recommendation at the end of the second carousel could be overtaken by the same recommendation but at the beginning of the third carousel, if it has a better position discount.

We repeated this procedure first optimizing NDCG, and then optimizing N2DCG. We consider a hypothetical interface with a total of 6 carousels, each composed of 10 items. The interface will initially show 3 carousels and 2 items. The user can display 1 additional item in a given carousel with each horizontal swipe and 1 new carousel with a vertical swipe. For this interface, we set $\alpha = \beta = 1$ and $\gamma = \lambda = 2$, in order to penalize more the swipes.

The resulting layouts are shown in Table 1. As we can see, the layouts have almost completely different orders of the chosen algorithms. For instance, optimizing N2DCG results in selecting SLIM as the first carousel, while the same algorithm was selected at the bottom of the layout that optimizes one-dimensional NDCG. UserKNN instead was the first algorithm when optimizing NDCG, but it is only the third carousel during N2DCG optimization.

Notice also how the 6 algorithms selected in both procedures are the same, only the order changes. Indeed, it is expected that NDCG and N2DCG will not produce completely different layouts but will differ the longer and more pronounced the effects of user actions become.

5 CONCLUSIONS

In this paper we have described a user interface with multiple carousels, typical of movie-on-demand and music streaming services, and based on its characteristics proposed an extended version

Optimizing NDCG	Optimizing N2DCG
UserKNN	SLIM
FunkSVD	FunkSVD
NMF	UserKNN
IALS	MF BPR
MF BPR	NMF
SLIM	IALS

Table 1: Layouts obtained optimizing NDCG and N2DCG.

of the widely used NDCG metric. The proposed formulation accounts for the known user behaviour of exploring the pages not one row at a time but focusing on the top-left corner and then navigating in both directions. The proposed formulation also allows to penalize correct recommendations that are only visible to the user after performing actions. Lastly, we show that the two metrics can lead to the selection of a different carousel layout. Future works include validating the proposed metric with user studies as well as applying it to select the optimal carousel layout, by defining which is the best carousel to put in a certain position or which is the best ordering of a given set of carousels. Also, further studies can be done on different gain and discount functions, similar to previous research works conducted on the one-dimensional DCG.

REFERENCES

- [1] Walid Bendada, Guillaume Salha, and Théo Bontempelli. 2020. Carousel Personalization in Music Streaming Apps with Contextual Bandits. In *RecSys 2020: Fourteenth ACM Conference on Recommender Systems, Virtual Event, Brazil, September 22-26, 2020*, Rodrygo L. T. Santos, Leandro Balby Marinho, Elizabeth M. Daly, Li Chen, Kim Falk, Noam Koenigstein, and Edleno Silva de Moura (Eds.). ACM, 420–425. <https://doi.org/10.1145/3383313.3412217>
- [2] Christopher J. C. Burges, Tal Shaked, Erin Renshaw, Ari Lazier, Matt Deeds, Nicole Hamilton, and Gregory N. Hullender. 2005. Learning to rank using gradient descent. In *Machine Learning, Proceedings of the Twenty-Second International Conference (ICML 2005), Bonn, Germany, August 7-11, 2005 (ACM International Conference Proceeding Series)*, Luc De Raedt and Stefan Wrobel (Eds.), Vol. 119. ACM, 89–96. <https://doi.org/10.1145/1102351.1102363>
- [3] Flavio Chierichetti, Ravi Kumar, and Prabhakar Raghavan. 2011. Optimizing two-dimensional search results presentation. In *Proceedings of the Fourth International Conference on Web Search and Web Data Mining, WSDM 2011, Hong Kong, China, February 9-12, 2011*, Irwin King, Wolfgang Nejdl, and Hang Li (Eds.). ACM, 257–266. <https://doi.org/10.1145/1935826.1935873>
- [4] Andrzej Cichocki and Anh Huy Phan. 2009. Fast Local Algorithms for Large Scale Nonnegative Matrix and Tensor Factorizations. *IEICE Trans. Fundam. Electron. Commun. Comput. Sci.* 92-A, 3 (2009), 708–721. <https://doi.org/10.1587/transfun.92.A.708>
- [5] Colin Cooper, Sang-Hyuk Lee, Tomasz Radzik, and Yiannis Siantos. 2014. Random walks in recommender systems: exact computation and simulations. In *23rd International World Wide Web Conference, WWW '14, Seoul, Republic of Korea, April 7-11, 2014, Companion Volume*, Chin-Wan Chung, Andrei Z. Broder, Kyuseok Shim, and Torsten Suel (Eds.). ACM, 811–816. <https://doi.org/10.1145/2567948.2579244>
- [6] Paolo Cremonesi, Yehuda Koren, and Roberto Turrin. 2010. Performance of recommender algorithms on top-n recommendation tasks. In *Proceedings of the 2010 ACM Conference on Recommender Systems, RecSys 2010, Barcelona, Spain, September 26-30, 2010*, Xavier Amatriain, Marc Torrens, Paul Resnick, and Markus Zanker (Eds.). ACM, 39–46. <https://doi.org/10.1145/1864708.1864721>
- [7] Ehtsham Elahi and Ashok Chandrashekar. 2020. Learning Representations of Hierarchical Slates in Collaborative Filtering. In *RecSys 2020: Fourteenth ACM Conference on Recommender Systems, Virtual Event, Brazil, September 22-26, 2020*, Rodrygo L. T. Santos, Leandro Balby Marinho, Elizabeth M. Daly, Li Chen, Kim Falk, Noam Koenigstein, and Edleno Silva de Moura (Eds.). ACM, 703–707. <https://doi.org/10.1145/3383313.3418484>
- [8] Maurizio Ferrari Dacrema, Simone Boglio, Paolo Cremonesi, and Dietmar Jannach. 2021. A Troubling Analysis of Reproducibility and Progress in Recommender Systems Research. *ACM Trans. Inf. Syst.* 39, 2, Article 20 (Jan. 2021), 49 pages. <https://doi.org/10.1145/3434185>

- [9] Alois Gruson, Praveen Chandar, Christophe Charbuillet, James McInerney, Samantha Hansen, Damien Tardieu, and Ben Carterette. 2019. Offline Evaluation to Make Decisions About Playlist Recommendation Algorithms. In *Proceedings of the Twelfth ACM International Conference on Web Search and Data Mining, WSDM 2019, Melbourne, VIC, Australia, February 11-15, 2019*, J. Shane Culpepper, Alistair Moffat, Paul N. Bennett, and Kristina Lerman (Eds.). ACM, 420–428. <https://doi.org/10.1145/3289600.3291027>
- [10] F. Maxwell Harper and Joseph A. Konstan. 2016. The MovieLens Datasets: History and Context. *ACM Trans. Interact. Intell. Syst.* 5, 4 (2016), 19:1–19:19. <https://doi.org/10.1145/2827872>
- [11] Jonathan L. Herlocker, Joseph A. Konstan, Loren G. Terveen, and John Riedl. 2004. Evaluating collaborative filtering recommender systems. *ACM Trans. Inf. Syst.* 22, 1 (2004), 5–53. <https://doi.org/10.1145/963770.963772>
- [12] Yifan Hu, Yehuda Koren, and Chris Volinsky. 2008. Collaborative Filtering for Implicit Feedback Datasets. In *Proceedings of the 8th IEEE International Conference on Data Mining (ICDM 2008), December 15-19, 2008, Pisa, Italy*. IEEE Computer Society, 263–272. <https://doi.org/10.1109/ICDM.2008.22>
- [13] Kalervo Järvelin and Jaana Kekäläinen. 2000. IR evaluation methods for retrieving highly relevant documents. In *SIGIR 2000: Proceedings of the 23rd Annual International ACM SIGIR Conference on Research and Development in Information Retrieval, July 24-28, 2000, Athens, Greece*, Emmanuel J. Yannakoudakis, Nicholas J. Belkin, Peter Ingwersen, and Mun-Kew Leong (Eds.). ACM, 41–48. <https://doi.org/10.1145/345508.345545>
- [14] Kalervo Järvelin and Jaana Kekäläinen. 2002. Cumulated gain-based evaluation of IR techniques. *ACM Trans. Inf. Syst.* 20, 4 (2002), 422–446. <https://doi.org/10.1145/582415.582418>
- [15] Kalervo Järvelin, Susan L. Price, Lois M. L. Delcambre, and Marianne Lykke Nielsen. 2008. Discounted Cumulated Gain Based Evaluation of Multiple-Query IR Sessions. In *Advances in Information Retrieval, 30th European Conference on IR Research, ECIR 2008, Glasgow, UK, March 30-April 3, 2008. Proceedings (Lecture Notes in Computer Science)*, Craig Macdonald, Iadh Ounis, Vassilis Plachouras, Ian Ruthven, and Ryan W. White (Eds.), Vol. 4956. Springer, 4–15. https://doi.org/10.1007/978-3-540-78646-7_4
- [16] Yvonne Kammerer and Peter Gerjets. 2010. How the interface design influences users' spontaneous trustworthiness evaluations of web search results: comparing a list and a grid interface. In *Proceedings of the 2010 Symposium on Eye-Tracking Research & Applications, ETRA 2010, Austin, Texas, USA, March 22-24, 2010*, Carlos Hitoshi Morimoto, Howell O. Istance, Aulikki Hyrskykari, and Qiang Ji (Eds.). ACM, 299–306. <https://doi.org/10.1145/1743666.1743736>
- [17] Evangelos Kanoulas and Javed A. Aslam. 2009. Empirical justification of the gain and discount function for nDCG. In *Proceedings of the 18th ACM Conference on Information and Knowledge Management, CIKM 2009, Hong Kong, China, November 2-6, 2009*, David Wai-Lok Cheung, Il-Yeol Song, Wesley W. Chu, Xiaohua Hu, and Jimmy J. Lin (Eds.). ACM, 611–620. <https://doi.org/10.1145/1645953.1646032>
- [18] Fernando Benjamín Pérez Maurera, Maurizio Ferrari Dacrema, Lorenzo Saulé, Mario Scriminaci, and Paolo Cremonesi. 2020. ContentWise Impressions: An Industrial Dataset with Impressions Included. In *CIKM '20: The 29th ACM International Conference on Information and Knowledge Management, Virtual Event, Ireland, October 19-23, 2020*, Mathieu d'Aquin, Stefan Dietze, Claudia Hauff, Edward Curry, and Philippe Cudré-Mauroux (Eds.). ACM, 3093–3100. <https://doi.org/10.1145/3340531.3412774>
- [19] Xia Ning and George Karypis. 2011. SLIM: Sparse linear methods for top-n recommender systems. In *Proceedings of the 11th IEEE International Conference on Data Mining (ICDM '11)*. 497–506.
- [20] Bibek Paudel, Fabian Christoffel, Chris Newell, and Abraham Bernstein. 2017. Updatable, Accurate, Diverse, and Scalable Recommendations for Interactive Applications. *ACM Trans. Interact. Intell. Syst.* 7, 1 (2017), 1:1–1:34. <https://doi.org/10.1145/2955101>
- [21] Steffen Rendle, Christoph Freudenthaler, Zeno Gantner, and Lars Schmidt-Thieme. 2009. BPR: Bayesian Personalized Ranking from Implicit Feedback. In *UAI 2009, Proceedings of the Twenty-Fifth Conference on Uncertainty in Artificial Intelligence, Montreal, QC, Canada, June 18-21, 2009*, Jeff A. Bilmes and Andrew Y. Ng (Eds.). AUAI Press, 452–461. https://dslpitt.org/uai/displayArticleDetails.jsp?mmnu=1&smnu=2&article_id=1630&proceeding_id=25
- [22] Mark Sanderson and W. Bruce Croft. 2012. The History of Information Retrieval Research. *Proc. IEEE* 100, Centennial-Issue (2012), 1444–1451. <https://doi.org/10.1109/JPROC.2012.2189916>
- [23] Badrul Sarwar, George Karypis, Joseph Konstan, and John Riedl. 2001. Item-based Collaborative Filtering Recommendation Algorithms. In *Proceedings of the 10th International Conference on World Wide Web (WWW '01)*. 285–295.
- [24] Badrul Munir Sarwar, George Karypis, Joseph A. Konstan, and John Riedl. 2001. Item-based collaborative filtering recommendation algorithms. In *Proceedings of the Tenth International World Wide Web Conference, WWW 10, Hong Kong, China, May 1-5, 2001*, Vincent Y. Shen, Nobuo Saito, Michael R. Lyu, and Mary Ellen Zurko (Eds.). ACM, 285–295. <https://doi.org/10.1145/371920.372071>
- [25] Harald Steck. 2019. Embarrassingly Shallow Autoencoders for Sparse Data. In *The World Wide Web Conference, WWW 2019, San Francisco, CA, USA, May 13-17, 2019*, Ling Liu, Ryan W. White, Amin Mantrach, Fabrizio Silvestri, Julian J. McAuley, Ricardo Baeza-Yates, and Leila Zia (Eds.). ACM, 3251–3257. <https://doi.org/10.1145/3308558.3313710>
- [26] Chao-Yuan Wu, Christopher V. Alvino, Alexander J. Smola, and Justin Basilico. 2016. Using Navigation to Improve Recommendations in Real-Time. In *Proceedings of the 10th ACM Conference on Recommender Systems, Boston, MA, USA, September 15-19, 2016*, Shilad Sen, Werner Geyer, Jill Freyne, and Pablo Castells (Eds.). ACM, 341–348. <https://doi.org/10.1145/2959100.2959174>
- [27] Qian Zhao, Shuo Chang, F. Maxwell Harper, and Joseph A. Konstan. 2016. Gaze Prediction for Recommender Systems. In *Proceedings of the 10th ACM Conference on Recommender Systems, Boston, MA, USA, September 15-19, 2016*, Shilad Sen, Werner Geyer, Jill Freyne, and Pablo Castells (Eds.). ACM, 131–138. <https://doi.org/10.1145/2959100.2959150>
- [28] Ke Zhou, Hongyuan Zha, Yi Chang, and Gui-Rong Xue. 2014. Learning the Gain Values and Discount Factors of Discounted Cumulative Gains. *IEEE Trans. Knowl. Data Eng.* 26, 2 (2014), 391–404. <https://doi.org/10.1109/TKDE.2012.252>