



A Logistic Factorization Model for Recommender Systems With Multinomial Responses

Yu Wang, Xuan Bi & Annie Qu

To cite this article: Yu Wang, Xuan Bi & Annie Qu (2020) A Logistic Factorization Model for Recommender Systems With Multinomial Responses, Journal of Computational and Graphical Statistics, 29:2, 396-404, DOI: [10.1080/10618600.2019.1665535](https://doi.org/10.1080/10618600.2019.1665535)

To link to this article: <https://doi.org/10.1080/10618600.2019.1665535>



View supplementary material [↗](#)



Accepted author version posted online: 11 Sep 2019.
Published online: 25 Oct 2019.



Submit your article to this journal [↗](#)



Article views: 149



View related articles [↗](#)



View Crossmark data [↗](#)



A Logistic Factorization Model for Recommender Systems With Multinomial Responses

Yu Wang^a, Xuan Bi^b, and Annie Qu^c

^aElectrical and Computer Engineering, Duke University, Durham, NC; ^bCarlson School of Management, University of Minnesota, Minneapolis, MN;

^cDepartment of Statistics, University of Illinois at Urbana-Champaign, Champaign, IL

ABSTRACT

In this article, we propose a two-way multinomial logistic model for recommender systems for categorical ratings. Specifically, we treat the possible ratings as mutually exclusive events, whose probability is determined by the latent factor of the users and the items through a two-way multinomial logistic function. The proposed method has a compatibility with categorical ratings and the advantage of incorporating both the covariate information and the latent factors of the users and items uniformly. We show numerically that the proposed method performs consistently better than five commonly used collaborative filtering methods, namely, the restricted singular value decomposition, the soft-impute matrix completion method, the regression-based latent factor models, the restricted Boltzmann machine, and the group-specific recommender system on various simulation setups and on MovieLens data. Supplementary materials for this article are available online.

ARTICLE HISTORY

Received April 2018

Revised August 2019

KEYWORDS

Cold-start problem;
Collaborative filter;
MovieLens

1. Introduction

The study of recommender systems (Aggarwal 2016) has flourished during the past decades (Resnick and Varian 1997; Ricci, Rokach, and Shapira 2011; Lu et al. 2015) stimulated by real-world challenges from Netflix (Bell and Koren 2007; Gomez-Uribe and Hunt 2015), Yahoo! Music (Koenigstein, Dror, and Koren 2011), and YouTube (Davidson et al. 2010), to name a few. Given a partial record of the ratings given by a set of users on a set of items, we would like to infer how a user will rate an item that has never been rated before based on the past behavior of the users and the rating history on the item (Ghosh et al. 1999; Adomavicius and Tuzhilin 2005; Kim et al. 2005; Aciar et al. 2006).

One of the major challenges in developing recommender systems that work in the real world is the sparsity of the rating matrix, as the number of the users and the items in a data is usually large and the average number of the ratings of a user or an item is relatively small. This leads to the situation where for many users and items in the testing set, there is no previous rating record in the training set to learn from. For example, in the MovieLens 10M data (Miller et al. 2003; Harper and Konstan 2015) collected by GroupLens Research (<http://grouplens.org/datasets/movielens>), 96% of the most recent ratings are from new users or on new items with no previous rating record (Bi et al. 2017). This phenomenon is called the “cold-start” problem. When tackling this problem, the covariate information on the new users and items, for example, the gender, sex, and age of a user and the features of an item can provide a relief.

Another feature of many datasets is that the ratings are categorical, taking values from a finite set of numbers. For example,

the ratings are integers from one to ten for IMDB (Oghina et al. 2012), and from one to five for the MovieLens data (Miller et al. 2003; Harper and Konstan 2015). This inspires using categorical data models (Agresti and Kateri 2011) to analyze the data.

Previously, two common approaches to predict a new rating are content-based filtering and collaborative filtering. For content-based filtering (Mooney and Roy 2000; Blanco-Fernandez et al. 2008; Iaquina et al. 2008; Lops, De Gemmis, and Semeraro 2011; Mirbakhsh and Ling 2015), a new rating of an item is generated by comparing its content with a user's profile. This method has the capability to handle “cold-start” problems with respect to items. In other words, the rating of a new item which has never been rated by any user before can be generated with relatively high accuracy (Levi et al. 2012). However, it cannot handle the “cold-start” problems with respect to users, since a past history is crucial to constructing a user's profile. In addition, sufficient domain knowledge or covariate information is usually required for content-based filtering to discriminate items the user likes from items the user does not like.

For collaborative filtering, the requirement for covariate information is relaxed by modeling it as latent factors which are to be inferred from existing ratings (Melville, Mooney, and Nagarajan 2002; Srebro, Alon, and Jaakkola 2005; Bobadilla et al. 2011; Cacheda et al. 2011; Ekstrand, Riedl, and Konstan 2011). Some of the most popular collaborative filtering approaches include restricted Boltzmann machines (Salakhutdinov, Mnih, and Hinton 2007) and various matrix factorization/completion methods: restricted singular value decomposition (Koren, Bell, and Volinsky 2009), the soft-impute matrix completion method (Mazumder, Hastie, and Tibshirani 2010), and the

regression-based latent factor models (Agarwal and Chen 2009). Practically, the performance of all these methods is comparable in most applications. However, they still cannot solve the “cold-start” problem effectively for users and the items simultaneously (Goldberg et al. 2000; Melville, Mooney, and Nagarajan 2002; Park et al. 2006; Nguyen, Denos, and Berrut 2007).

In this work, we propose a logistic-regression based collective filtering method that consistently unifies content-based filtering and collaborative filtering. Previously, for existing collaborative filtering methods, the covariate information of the users and items is either not easy to incorporate, or is used solely as an “offset” to the ratings (Agarwal and Chen 2009; Cron, Zhang, and Agarwal 2014), or to group users and the items (Bi et al. 2017). In our method, both the covariate information and the latent factors of the users and the items are uniformly incorporated in the model—they are combined together and contribute equally to the preference of a user giving or an item receiving a specific rating via a multinomial logistic method. This promises the proposed method a better performance in solving the cold-start problem.

In addition, in our method, we use latent factors to generate the *probability mass functions* of the ratings. This is rather different from the matrix factorization/completion based collaborative filtering methods (Agarwal and Chen 2009; Koren, Bell, and Volinsky 2009; Stern, Herbrich, and Graepel 2009; Mazumder, Hastie, and Tibshirani 2010; Zhu, Shen, and Ye 2016; Bi et al. 2017), where the latent factors are used to generate the values of the ratings directly.

Admittedly, the matrix-factorization/completion-methods often boast nice theoretical guarantees on the optimality of the solutions in the sense of minimizing the root mean square error of the predicted ratings with respect to the true ratings when the ratings are generated by the proposed models; however, most of these results are only valid when the ratings take values in real numbers. In many applications, the ratings are categorical (Miller et al. 2003; Oghina et al. 2012; Harper and Konstan 2015). In particular, the ratings may be binary in certain circumstances—for example, on a photo or video-sharing website like Youtube or Instagram, a photo or video is either liked or disliked by a user. In these cases, our logistic-regression-based collective filtering method is a more natural and suitable choice than the matrix-factorization/completion-methods, especially when the ratings are binary.

Compared to the restricted Boltzmann machine (Salakhutdinov, Mnih, and Hinton 2007) that takes a similar approach of treating possible ratings as mutually exclusive events whose probability distribution is generated by some hidden variable, our approach has two advantages. First, in our method, the latent factors are associated with a specific user or item instead of being shared within the whole model. Thus they have clearer interpretation and are more capable of distinguishing the different preferences of the users and the items. In addition, it is much more difficult to incorporate covariate information associated to a user or an item in the restricted Boltzmann machine than in our model.

The rest of the article is organized as follows. In Section 2, we provide the notations and background that will be used in the rest of the article. In Section 3, we introduce the two-way multinomial logistic model with or without covariate informa-

tion and give the algorithm to fit the model into a given set of ratings. In Section 4, we implement the proposed method on simulated binary rating data under different missing rates and under different data sizes, simulated multicategorical rating data under different missing rates and in the cold-start problem, and compare the results to the restricted singular value decomposition (Koren, Bell, and Volinsky 2009), the soft-impute matrix completion method (Mazumder, Hastie, and Tibshirani 2010), the regression-based latent factor models (Agarwal and Chen 2009), the restricted Boltzmann machine (Salakhutdinov, Mnih, and Hinton 2007), and the group-specific recommender system (Bi et al. 2017). In Section 5, we implement the proposed method on MovieLens data and compared the results to the same existing methods. Finally, we conclude this work in Section 6.

2. Notations and Background

In the rest of the article, we denote the set of real numbers and natural numbers by $\mathbb{R}$ and $\mathbb{N}$, respectively. For $n \in \mathbb{N}$, let $[n] = \{1, \dots, n\}$. Let $|\cdot|$ be the cardinality of a set. The indicator function is denoted by $\mathbf{1}_{\{\cdot\}}$. Let e_n be the n th natural basis in the Euclidean space, namely, all the entries in e_n are 0 except the n th entry being 1.

Consider m users rating n items with possible ratings from $\{1, \dots, r\}$. The records, which may not be complete, can be represented by an $m \times n$ utility matrix $\mathbf{R}$, whose (i, j) entry $\mathbf{R}_{ij}$ gives the score of item j by user i , if it exists; and 0, otherwise. Let $\Omega = \{(i, j) \in [m] \times [n] \mid \mathbf{R}_{ij} \neq 0\}$ be the set of entries where the rating exists, $\Omega_i = \{j \in [n] \mid \mathbf{R}_{ij} \neq 0\}$ be the set of items rated by user i , and $\Omega_j = \{i \in [m] \mid \mathbf{R}_{ij} \neq 0\}$ be the set of users rating item j . By slightly modifying the notation, for $R \in \mathbb{R}^{m \times n}$, let $\Omega(R)$ be

$$\Omega(\mathbf{R})_{ij} = \begin{cases} \mathbf{R}_{ij}, & \text{if } (i, j) \in \Omega \\ 0, & \text{otherwise.} \end{cases}$$

The matrix factorization/completion based collaborative filtering methods (Agarwal and Chen 2009; Koren, Bell, and Volinsky 2009; Mazumder, Hastie, and Tibshirani 2010; Zhu, Shen, and Ye 2016; Bi et al. 2017) are based on decomposing the utility matrix $\mathbf{R}$ approximately as

$$\Omega(\mathbf{R}) \approx \Omega(\mathbf{P}\mathbf{Q}^T), \quad (1)$$

where $\mathbf{P}$ is an $m \times l$ user preference matrix, $\mathbf{Q}$ an $n \times l$ item preference matrix, and l is the given number of latent factors. Different proximity criteria lead to different methods. For example, the residual after a baseline fit, similar to ANOVA, is suggested in Koren (2010) and Mazumder, Hastie, and Tibshirani (2010); and ℓ_0 and ℓ_1 are suggested in Zhu, Shen, and Ye (2016). To estimate $\mathbf{P}$ and $\mathbf{Q}$ numerically, it is common to use alternating least square methods or the gradient descent methods (Koren, Bell, and Volinsky 2009). Finally, the predicted value of $\mathbf{R}_{ij}$ for $(i, j) \in \Omega$ is given by $\mathbf{P}_{ij}$.

On the other hand, the restricted Boltzmann machine method (Salakhutdinov, Mnih, and Hinton 2007) is based on the belief that the preference of each user is represented by a set of binary or Gaussian hidden units $(h_1, \dots, h_F)$, whose values are different among the users. And a user rates the movie i by

$\mathbf{R}_i = k$ with probability

$$\mathbb{P}[\mathbf{R}_i = k] \propto e^{\mathbf{b}_{ik} + \sum_{f=1}^F h_f \mathbf{W}_{ikf}}, \quad (2)$$

for each possible rating $k = 1, \dots, r$. Given a partial observation of the ratings, the parameters $\mathbf{b}_{ik}$, $\mathbf{W}_{ikf}$ and the distribution of the value of the hidden units h_f can be estimated by MCMC or “contrastive divergence” (Salakhutdinov, Mnih, and Hinton 2007). Then the probability distribution of the missing ratings can be derived using (2) again, and the expectation can be taken as the predicted ratings.

3. Method

3.1. The Proposed Model

In this model, each user i has latent factors $\mathbf{u}_i = (\mathbf{u}_{i1}, \dots, \mathbf{u}_{ir})$, and each item j has latent factors $\mathbf{v}_j = (\mathbf{v}_{j1}, \dots, \mathbf{v}_{jr})$, where $\mathbf{u}_{i1}, \dots, \mathbf{u}_{ir}, \mathbf{v}_{j1}, \dots, \mathbf{v}_{jr} \in \mathbb{R}^l$. In addition, there are latent factors $\mathbf{c}_k$ and $\mathbf{d}_k$ that describe the overall preferences of the users and the items, respectively. We note that each $\mathbf{u}_{ik}$, $\mathbf{v}_{jk}$, $\mathbf{c}_k$, and $\mathbf{d}_k$ are vectors of the same dimension l . The rating of item j by user i is assumed to follow the probability mass function

$$\mathbb{P}[\mathbf{R}_{ij} = k] = \frac{e^{(\mathbf{u}_{ik} + \mathbf{c}_k)^T (\mathbf{v}_{jk} + \mathbf{d}_k)}}{\sum_{l=1}^r e^{(\mathbf{u}_{il} + \mathbf{c}_l)^T (\mathbf{v}_{jl} + \mathbf{d}_l)}}, \quad k \in [r]. \quad (3)$$

For user i with latent factors $(\mathbf{u}_{i1} + \mathbf{c}_1, \dots, \mathbf{u}_{ir} + \mathbf{c}_r)$, the probability that it rates k on the item j is only proportional to the exponential of the k th coefficient of the item; and the same for the items.

Therefore, this model can be viewed as a two-way multinomial logistic model. The model (3) is similar to the restricted Boltzmann machine (2) in the sense that the possible ratings are treated as mutually exclusive events, whose probability distribution is determined by the latent factors. Therefore, it is a more natural and suitable choice than the matrix-factorization/completion-methods for categorical ratings.

When the covariate information of the users and the items are available, we can incorporate them into the model by augmentation. If each user i has m_{CI} covariate information $(\boldsymbol{\alpha}_1, \dots, \boldsymbol{\alpha}_{m_{\text{CI}}})$, and each item j has n_{CI} covariate information $(\boldsymbol{\beta}_1, \dots, \boldsymbol{\beta}_{n_{\text{CI}}})$, then we can augment the latent factors by

$$\tilde{\mathbf{u}}_{ik} = (\mathbf{u}_{ik}, \boldsymbol{\alpha}_{i_{\text{CI}}}, \mathbf{b}_{i_{\text{CI}}k}), \quad \tilde{\mathbf{v}}_{jk} = (\mathbf{v}_{jk}, \mathbf{a}_{j_{\text{CI}}k}, \boldsymbol{\beta}_{j_{\text{CI}}k}), \quad (4)$$

where $\mathbf{a}_{i_{\text{CI}}}$, $\mathbf{b}_{j_{\text{CI}}}$ are the latent factors of proper dimensions associated with the covariate information. In this case, we have

$$p_{ijk} = \mathbb{P}[\mathbf{R}_{ij} = k] = \frac{e^{(\tilde{\mathbf{u}}_{ik} + \mathbf{c}_k)^T (\tilde{\mathbf{v}}_{jk} + \mathbf{d}_k)}}{\sum_{l=1}^r e^{(\tilde{\mathbf{u}}_{il} + \mathbf{c}_l)^T (\tilde{\mathbf{v}}_{jl} + \mathbf{d}_l)}}, \quad k \in [r]. \quad (5)$$

Since in this model the latent factors are associated with a specific user or item instead of being shared within the whole model, thus they are more capable of telling the different preferences of the users and the items. More importantly, this makes it possible to uniformly incorporate covariate information, unlike the restricted Boltzmann machine.

Given the augmented latent factors $\tilde{\mathbf{u}}_{ik}$, $\tilde{\mathbf{v}}_{jk}$, $\mathbf{c}_k$, and $\mathbf{d}_k$, the predicted ratings are generated depending on the evaluation criteria (Gunawardana and Shani 2015). Similar to previous works (Salakhutdinov, Mnih, and Hinton 2007; Agarwal and Chen 2009; Mazumder, Hastie, and Tibshirani 2010; Bi et al. 2017), we use the root mean square error (RMSE) to evaluate the predictions and consider the following objective function

$$\mathcal{S} = \sum_{(i,j) \in \Omega} (\bar{\mathbf{R}}_{ij} - \mathbf{R}_{ij})^2, \quad (6)$$

where the predicted ratings are the expectation of all possible ratings

$$\bar{\mathbf{R}}_{ij} = \sum_{k=1}^r k p_{ijk}. \quad (7)$$

Compared to the maximum likelihood estimation (MLE) approach, the RMSE approach is more robust to irregular data, such as in the examples where binary covariates can completely separate binary responses. In addition, the RMSE approach is more computationally effective in handling missing data.

3.2. Estimating the Latent Factors

Given a record of ratings $\mathbf{R}_{ij}$ for $(i, j) \in \Omega$, the latent factors $\mathbf{u}_{ik}$, $\mathbf{v}_{jk}$, $\mathbf{c}_k$, and $\mathbf{d}_k$ can be estimated by minimizing the mis-prediction rate or the root mean square error by applying the gradient descent method iteratively on the latent factors. The gradients of the objective function (7) with respect to the latent factors $\mathbf{u}_{ik}$, $\mathbf{v}_{jk}$, $\mathbf{c}_k$, and $\mathbf{d}_k$ are

$$\begin{aligned} \frac{\partial \mathcal{S}}{\partial \tilde{\mathbf{u}}_{ik}} &= \sum_{j \in \Omega_i} P_{ijk} (\tilde{\mathbf{v}}_{jk} + \mathbf{d}_k), & \frac{\partial \mathcal{S}}{\partial \tilde{\mathbf{v}}_{jk}} &= \sum_{i \in \Omega_j} P_{ijk} (\tilde{\mathbf{u}}_{ik} + \mathbf{c}_k), \\ \frac{\partial \mathcal{S}}{\partial \mathbf{c}_k} &= \sum_{(i,j) \in \Omega} P_{ijk} (\tilde{\mathbf{v}}_{jk} + \mathbf{d}_k), & \frac{\partial \mathcal{S}}{\partial \mathbf{d}_k} &= \sum_{(i,j) \in \Omega} P_{ijk} (\tilde{\mathbf{u}}_{ik} + \mathbf{c}_k), \end{aligned} \quad (8)$$

where

$$P_{ijk} = 2(\bar{\mathbf{R}}_{ij} - \mathbf{R}_{ij})(k - \bar{\mathbf{R}}_{ij})p_{ijk}. \quad (9)$$

To implement the gradient descent, we set the initial values for the latent factors as

$$\begin{aligned} \mathbf{u}_{ik}, \mathbf{v}_{jk} &\sim N(0, \eta I_D), \quad \mathbf{a}_{i_{\text{CI}}k}, \mathbf{b}_{j_{\text{CI}}k} = 0, \\ \mathbf{c}_k = \mathbf{d}_k &= \frac{1}{2} \log \left(\sum_{(i,j) \in \Omega} \mathbf{1}_{\{\mathbf{R}_{ij}=k\}} / l \right), \quad k \in [r], \end{aligned} \quad (10)$$

where l is the dimension of each $\mathbf{u}_{ik}$, $\mathbf{v}_{jk}$, $\mathbf{c}_k$, and $\mathbf{d}_k$, and $\eta \ll \min \mathbf{c}_k$. In the simulations, we choose $\eta = \min \mathbf{c}_k / 10$. The choice of $\mathbf{c}_k$ and $\mathbf{d}_k$ is to ensure that when $\mathbf{u}_{ik}$ and $\mathbf{v}_{jk}$ are small, namely, the users and items have no strong preference, and the probability mass function of the predicted rating is approximated by the empirical distribution of all the existing ratings.

3.3. Implementation Issues

We use the stochastic gradient descent method with decreasing step sizes to fit the model with existing data (Spall 2005). The purpose is to ensure accurate small-step searching near the local minimum, while allowing for occasional big steps for a faster convergence. The step sizes of the gradient descent process for the latent factors $\mathbf{u}_{ik}$, $\mathbf{v}_{jk}$, $\mathbf{c}_k$, and $\mathbf{d}_k$ are randomly generated, respectively, by

$$\begin{aligned}\delta_{\mathbf{u}_{ik}} &= \Delta_K(\partial_{\mathbf{u}_{ik}} \mathcal{S}(t)), & \delta_{\mathbf{v}_{jk}} &= \Delta_K(\partial_{\mathbf{v}_{jk}} \mathcal{S}(t)), \\ \delta_{\mathbf{c}_k} &= \Delta_K(\partial_{\mathbf{c}_k} \mathcal{S}(t)), & \delta_{\mathbf{d}_k} &= \Delta_K(\partial_{\mathbf{d}_k} \mathcal{S}(t)), \\ \delta_{\mathbf{a}_{iClk}} &= \Delta_K(\partial_{\mathbf{a}_{iClk}} \mathcal{S}(t)), & \delta_{\mathbf{b}_{jClk}} &= \Delta_K(\partial_{\mathbf{b}_{jClk}} \mathcal{S}(t)),\end{aligned}\quad (11)$$

where the $\Delta_K(x)$ are sampled from the exponential distribution $\text{EXP}(\max\{x, K\})$ independently with a rate of $\max\{x, K\}$ and K controlling the step size. Consequently, the latent factors update by

$$\begin{aligned}\mathbf{u}_{ik} &\leftarrow \mathbf{u}_{ik} - \delta_{\mathbf{u}_{ik}} \partial_{\mathbf{u}_{ik}} \mathcal{S}(t), & \mathbf{v}_{jk} &\leftarrow \mathbf{v}_{jk} - \delta_{\mathbf{v}_{jk}} \partial_{\mathbf{v}_{jk}} \mathcal{S}(t), \\ \mathbf{c}_k &\leftarrow \mathbf{c}_k - \delta_{\mathbf{c}_k} \partial_{\mathbf{c}_k} \mathcal{S}(t), & \mathbf{d}_k &\leftarrow \mathbf{d}_k - \delta_{\mathbf{d}_k} \partial_{\mathbf{d}_k} \mathcal{S}(t), \\ \mathbf{a}_{iClk} &\leftarrow \mathbf{a}_{iClk} - \delta_{\mathbf{a}_{iClk}} \partial_{\mathbf{a}_{iClk}} \mathcal{S}(t), & \mathbf{b}_{jClk} &\leftarrow \mathbf{b}_{jClk} - \delta_{\mathbf{b}_{jClk}} \partial_{\mathbf{b}_{jClk}} \mathcal{S}(t).\end{aligned}\quad (12)$$

The convergence in the gradient descent is checked by comparing the relative change in the objective function within T steps. Specifically, if the objective function changes less than ϵ , relatively, in T steps, then we stop the algorithm. The above statements are concluded in Algorithm 1.

Algorithm 1 Infer the latent factors by minimizing the misprediction rate or the root mean square error

-
- 1: **Input:** rating matrix $\mathbf{R}$, convergence threshold T, ϵ , parameter K , $t = 1$
 - 2: Initialize the latent factors $\tilde{\mathbf{u}}_{ik}, \tilde{\mathbf{v}}_{jk}, \mathbf{c}_k, \mathbf{d}_k$ by (10)
 - 3: Compute the value of the objective function $\mathcal{S}(t)$ by (7)
 - 4: **repeat**
 - 5: Compute the new latent factors by gradient descent by (12)
 - 6: Compute the new value of the objective function $\mathcal{S}'$ by (7)
 - 7: **if** $\mathcal{S}' < \mathcal{S}(t)$ **then**
 - 8: Update the latent factors $\tilde{\mathbf{u}}_{ik}, \tilde{\mathbf{v}}_{jk}, \mathbf{c}_k, \mathbf{d}_k$ and set $\mathcal{S}(t + 1) = \mathcal{S}'$
 - 9: **else**
 - 10: Set $\mathcal{S}(t + 1) = \mathcal{S}(t)$ and **continue**
 - 11: **end if**
 - 12: $t = t + 1$
 - 13: **until** $t > T$ and $\frac{|\mathcal{S}(t) - \mathcal{S}(t-T)|}{\mathcal{S}(t-T)} < \epsilon$.
-

4. Simulation

In this section, we implement the two-way multinomial logistic model with or without covariate information and compare it to other existing collaborative filtering methods, namely the restricted singular value decomposition (Koren, Bell, and Volinsky 2009), the soft-impute matrix completion method (Mazumder, Hastie, and Tibshirani 2010), the regression-based latent factor models (Agarwal and Chen 2009), the restricted

Boltzmann machine (Salakhutdinov, Mnih, and Hinton 2007), and the group-specific recommender system (Bi et al. 2017) on the simulated binary rating data in Section 4.1 and on the simulated multicategorical rating data in Section 4.2 to study the performance of our method under different missing rates and in the cold-start problem.

For the restricted singular value decomposition (Koren, Bell, and Volinsky 2009), we use the R package. For the soft-impute matrix completion method (Mazumder, Hastie, and Tibshirani 2010), we apply the R package “softImpute.” For the regression-based latent factor models (Agarwal and Chen 2009), we use the default 10 iterations. For the group-specific recommender system (Bi et al. 2017), we use the Matlab code available at <https://sites.google.com/site/xuanbigts/software>. For the restricted Boltzmann machine (Salakhutdinov, Mnih, and Hinton 2007), we use the python package TensorFlow available at <https://github.com/tensorflow/tensorflow>. The two-way multinomial logistic model with or without covariate information is implemented in Matlab available at <https://bitbucket.org/yuwang0531/tml>.

4.1. Binary Classification

In this section, we study the performance of the proposed method on the simulated binary rating data. This is the kind of problem that the matrix factorization/completion based methods are not specialized for. The simulated binary rating data are generated by the two-way multinomial logistic model with covariate information introduced in Section 3 with the number of category $r = 2$. Due to similarity in the performance of the competing models in fitting real data (the best one is only 10% better than the worst one), and the relatively high missing rate, the trends of the results on simulated data are the same no matter which model is used to generate the ratings. (We have performed an additional simulation with 5-category ratings generated by 1500 users and 1500 items under the RSVD model in the supplementary materials. The trends are similar to the results derived from the simulated data generated by the proposed model.) Using the simulated data, we show the advantage of our method compared to the other methods under different missing rates in Section 4.1.1 and under different data sizes in Section 4.1.2.

4.1.1. Missing Rate

To study the influence of the missing rate, we generate the ratings for 300 users and 300 items under the missing rates 0.7, 0.8, 0.9, 0.95. The dimension of latent factors is $l = 3$. The numbers of covariate information for both the items and the users are $m_{Cl} = 2$ and $n_{Cl} = 2$, respectively. The latent factors $\mathbf{u}_{ik}, \mathbf{v}_{jk}, \mathbf{c}_k, \mathbf{d}_k$ are independently generated from the Gaussian distribution $N(0, 0.1)$. The covariate information for the users α_{iCl} and the times β_{iCl} is independently generated from the Bernoulli distribution $B(0.5)$ in $\{0, 1\}$. The corresponding latent factors $\mathbf{a}_{iClk}$ and $\mathbf{b}_{jClk}$ are independently generated from the Gaussian distribution $N(0, 0.1)$. To generate the ratings at the desired missing rate, we first generate all possible $m \times n$ ratings between the users and the items and then pick out a given portion of the ratings. Also, all the ratings are randomly permuted to simulate random log-ins of users in real applications and divided

by 60% for training, 15% for tuning, and 25% for testing, to mimic the setting of MovieLens data (see Section 5).

For the regularized singular value decomposition method (Koren, Bell, and Volinsky 2009), we set the tuning parameter $\lambda = 3$. For the regression-based latent factor model (Agarwal and Chen 2009), we use the default of 10 iterations. For the soft-impute method (Mazumder, Hastie, and Tibshirani 2010), we choose the default $\lambda = 0, K = 4$ to achieve convergence for the local minimum. For the group-specific method (Bi et al. 2017), we choose $K = 3$. For the restricted Boltzmann machine (Salakhutdinov, Mnih, and Hinton 2007), 100 hidden units are used. The two-way multinomial logistic model with or without covariate information is implemented with dimension of latent factors $l = 3$ for the best trade-off between accuracy and computational cost.

The simulation results for the above methods are shown in Table 1, where the mean is calculated from 5000 runs, and the standard error of the mean is less than 0.002. It shows that the advantage of the two-way logistic model with covariate information is generally more significant under high missing rate. For example, under the missing rate of 95%, the root mean square error of the two-way logistic model with covariate information is 4.3% smaller than the same model without covariate information, 10.5% smaller than that of the regularized singular value decomposition method (Koren, Bell, and Volinsky 2009), 7.4% smaller than the restricted Boltzmann machine (Salakhutdinov, Mnih, and Hinton 2007), 7.0% smaller than the regression-based latent factor model (Agarwal and Chen 2009), 6.0% smaller than the soft-impute method (Mazumder, Hastie, and Tibshirani 2010), and 5.5% smaller than the group-specific recommender system (Bi et al. 2017). This indicates that the covariate information is more important under high missing rate and the two-way logistic model utilizes the covariate information more effectively.

4.1.2. Data Size

To study the influence of the data size, we generate ratings for 300 users and 300 items, and for 1500 users and 1500 items, under the missing rates 0.8, 0.95 from the model described in Section 3. The dimension of latent factors is $l = 3$. The numbers

of covariate information for both the items and the users are $m_{CI} = 2$ and $n_{CI} = 2$, respectively. The latent factors for the users and items and ratings in the training, tuning and testing sets are generated in the same way as in Section 4.1.1. The parameter settings of the regularized singular value decomposition method (Koren, Bell, and Volinsky 2009), the regression-based latent factor model (Agarwal and Chen 2009), the soft-impute method (Mazumder, Hastie, and Tibshirani 2010), the group-specific method (Bi et al. 2017), the restricted Boltzmann machine (Salakhutdinov, Mnih, and Hinton 2007), and two-way multinomial logistic model with or without covariate information are also the same as given in Section 4.1.1.

The simulation results for the above methods are shown in Table 2, where the mean is calculated from 5000 runs, and the standard error of the mean is less than 0.002. It shows that the advantage of the two-way logistic model with covariate information is consistent for different data sizes. For all the methods, the performance is better on the larger data than in the smaller data. This is because under the same missing rate, the number of ratings for each user and each item increases as the data size increases, thus it is easier to make predictions on the larger data size.

Table 2 also shows that the two-way logistic model with covariate information is less sensitive to data size change than other methods. Namely, this method is better at handling smaller data than the other methods. For example, under a missing rate of 95%, the increase in the root mean square error after changing the data size from 1500×1500 to 300×300 for the two-way logistic model with covariate information is 3.80%, while the increments are 6.95%, 5.51%, 5.84%, 5.84%, 6.79%, 6.27% for the regularized singular value decomposition method, the restricted Boltzmann machine, the regression-based latent factor model, the soft-impute method, the group-specific recommender system, and the two-way logistic model without covariate information, respectively.

4.2. Multicategorical Recommendation With Cold-Start

In this section, we study the performance of the proposed method on multicategorical data. This is the most common

Table 1. Comparison of the root mean square error of the methods under different missing rates on the simulated binary rating data with 300 users and 300 items.

Missing rate	RSVD	RBM	RB	SI	GS	TML	TMLCI
0.7	0.3784	0.3812	0.3742	0.3735	0.3624	0.3543	0.3484
0.8	0.3984	0.3904	0.3843	0.3793	0.3745	0.3635	0.3584
0.9	0.4064	0.3983	0.3997	0.3913	0.3874	0.3793	0.3643
0.95	0.4184	0.4043	0.4026	0.3984	0.3965	0.3913	0.3745

NOTE: RSVD, RBM, RB, SI, GS, TML, and TMLCI stand for the regularized singular value decomposition method (Koren, Bell, and Volinsky 2009), the restricted Boltzmann machine (Salakhutdinov, Mnih, and Hinton 2007), the regression-based latent factor model (Agarwal and Chen 2009), the soft-impute method (Mazumder, Hastie, and Tibshirani 2010), the group-specific recommender system (Bi et al. 2017), the proposed two-way multinomial logistic model, and the proposed two-way multinomial logistic model with covariate information, respectively.

Table 2. Comparison of the methods under different sample sizes on the simulated binary rating data.

Missing rate	Data size	RSVD	RBM	RB	SI	GS	TML	TMLCI
0.8	300×300	0.3984	0.3904	0.3843	0.3793	0.3745	0.3635	0.3584
	1500×1500	0.3774	0.3713	0.3647	0.3573	0.3532	0.3484	0.3454
0.95	300×300	0.4184	0.4043	0.4026	0.3984	0.3965	0.3913	0.3795
	1500×1500	0.3912	0.3832	0.3804	0.3764	0.3713	0.3682	0.3656

setup for recommender systems in real applications. The simulated multicategorical rating data are generated by the two-way multinomial logistic model with covariate information introduced in Section 3, with five categories $\{1, 2, 3, 4, 5\}$ in the same fashion as the MovieLens data. Using simulated data, we show the advantage of our method compared to the other existing methods, under different missing rates again in Section 4.2.1, and in the cold-start problem (Section 4.2.2).

4.2.1. Missing Rate

We generate ratings for 1500 users and 1500 items under the missing rates 0.7, 0.8, 0.9, 0.95. The dimension of latent factors is $l = 2$. The numbers of covariates for both the items and the users are $m_{CI} = 1$ and $n_{CI} = 1$, respectively. The latent factors $\mathbf{u}_{ik}, \mathbf{v}_{jk}, \mathbf{c}_k, \mathbf{d}_k$ are independently generated from the Gaussian distribution $N(0, 0.1)$. The covariate information for the users α_{iCI} and the times β_{iCI} is independently generated from the Bernoulli distribution $B(0.5)$ in $\{0, 1\}$. The corresponding latent factors $\mathbf{a}_{iCIk}$ and $\mathbf{b}_{iCIk}$ are independently generated from the Gaussian distribution $N(0, 0.1)$. To generate the ratings at the desired missing rate, we first generate all possible $m \times n$ ratings between the users and the items and then pick out a given portion of the ratings. All the ratings are randomly permuted to simulate random log-ins of the users in the real case, and divided by 60% for training, 15% for tuning and 25% for testing.

The parameters for the competing methods are tuned for their best performance, as follows. Specifically, we set $K = 4$ and $\lambda = 6$ for the regularized singular value decomposition method, choose $\lambda = 0$ and 10 iterations for the regression-based latent factor model, select $K = 4$ for the soft-impute method, and choose 100 hidden units for the restricted Boltzmann machine. The two-way multinomial logistic model with or without covariate information is implemented with dimension of latent factors $l = 2$ for the objective function (6).

The simulation results for the above methods are shown in Table 3, where the mean is calculated from 10,000 runs, and the standard error of the mean is less than 0.002. Similar to the trend shown in Table 1, the advantage of the two-way logistic model with covariate information is generally more significant under high missing rates. For example, under the missing rate of 95%, the root mean square error of the two-way logistic model with covariate information is 2.3% smaller than the same model without covariate information, 12.9% smaller than that of the regularized singular value decomposition method,

11.5% smaller than the restricted Boltzmann machine, 14.8% the regression-based latent factor model, 11.0% smaller than the soft-impute method, and 1.8% smaller than the group-specific recommender system.

4.2.2. Cold-Start Problem

In this simulation study, we consider the “cold-start” problem where the testing set contains a large portion of new users and new items that have not appeared in the training set. We generate ratings for 1500 users and 1500 items under a missing rate of 0.99. The dimension of latent factors is $l = 2$. The numbers of covariates for both the items and the users are $m_{CI} = 1$ and $n_{CI} = 1$, respectively. The latent factors for the users and the items are generated in the same way as in Section 4.2.1. The parameter settings of the regularized singular value decomposition method, the regression-based latent factor model, the soft-impute method, the group-specific method, the restricted Boltzmann machine, and two-way multinomial logistic model with or without covariant are also the same as given in Section 4.2.1.

The ratings are generated at a missing rate of 0.99 and randomly permuted to simulate random log-ins of the users in the real case. They are divided by 75% for training, 10% for tuning, and 15% for testing. To create a cold-start setup, we then exchange the ratings in the testing set with the ratings in the training and tuning sets, such that 50% of the ratings in the testing set are either from new users or for new items.

The simulation results for the above methods are shown in Table 4, where the mean is calculated from 10,000 runs, and the standard error of the mean is less than 0.002. It shows that the two-way logistic model with or without covariate information outperforms the other methods in the cold-start problem. Specifically, on the whole testing set, the root mean square error of the two-way logistic model with covariate information is 1.3% smaller than the same model without covariate information, 9.8% smaller than that of the regularized singular value decomposition method, 9.3% smaller than the restricted Boltzmann machine, 12.0% the regression-based latent factor model, 8.8% smaller than the soft-impute method, and 2.2% smaller than the group-specific recommender system. For the new ratings, the root mean square error of the two-way logistic model with covariate information is 1.7% smaller than the same model without covariate information, 11.1% smaller than that of the regularized singular value decomposition method,

Table 3. Comparison of the methods under different missing rates on the simulated multicategory rating data with 1500 users and 1500 items.

Missing rate	RSVD	RBM	RB	SI	GS	TML	TMLCI
0.7	1.0523	1.0341	1.0764	1.0352	0.9724	0.9754	0.9514
0.8	1.0584	1.0404	1.0954	1.0567	0.9745	0.9721	0.9537
0.9	1.0745	1.0632	1.1177	1.0632	0.9798	0.9823	0.9601
0.95	1.1084	1.0903	1.1326	1.0844	0.9834	0.9878	0.9654

Table 4. Comparison of the methods in the cold-start problem on the simulated five-category rating data.

	RSVD	RBM	RB	SI	GS	TML	TMLCI
Old	0.9761	0.9428	0.9817	0.9739	0.9221	0.9175	0.9013
New	1.2033	1.2204	1.2494	1.1834	1.0923	1.0799	1.0692
Average	1.0956	1.0905	1.1236	1.0837	1.0108	1.0020	0.9888

12.4% smaller than the restricted Boltzmann machine, 14.4% the regression-based latent factor model, 9.7% smaller than the soft-impute method, and 2.3% smaller than the group-specific recommender system.

5. MovieLens Data

We apply the proposed method to the MovieLens 1M and 10M data (Miller et al. 2003; Harper and Konstan 2015), collected by GroupLens Research (<http://grouplens.org/datasets/movielens>). The MovieLens 1M data contain 1,000,209 integer ratings of 3883 movies by 6040 users ranging from $\{1, 2, \dots, 5\}$, at a missing rate of 96%. In addition, it provides demographic information including the age, gender, occupation, and zip code of the users, and the genres of the movies. In the MovieLens 10M data, 10,000,054 ratings are collected from 71,567 users over 10,681 items, at a missing rate of 99%. For the 10M data, the user's demographic information is not available.

The categorical covariate information of the users and the items is binary-coded. To avoid using too many variables, we divide users' ages by 0–9, 10–19, ..., 60–69 into seven groups, considering that users of similar ages tend to have similar preference for movies. The users' zip codes are divided into 10 groups by 0–9999, ..., 90,000–99,999, since users of similar zip codes tend to live closer geographically. The users' gender is either male or female, hence encoded by $\{0, 1\}$. In addition, each user has one of 21 possible occupations. Consequently, the users' covariate information is encoded into 39 binary variables. Each item is labeled by a subset of 18 possible genres; therefore, the items' covariate information is encoded by 18 binary variables.

The simulation results for the proposed methods and the other existing methods are shown in Table 5. The data are divided by 60% for training, 15% for tuning, and 25% for testing, sorted by timestamps. For the proposed two-way multinomial logistic model with or without covariate information, we apply the loss function (7) and set the parameter $l = 5$ to achieve the best performance with the smallest value for l , after searching over $l = 1, 2, \dots, 10$. The parameters for competing models are tuned for their best performance as follows. For the RSVD, we use $K = 4$ and $\lambda = 8$ for the 1M data, and $K = 4$ and $\lambda = 6$ for the 10M data. For the regression-based latent factor model, we let $K = 1$ for both the 1M and 10M data and take 25 and 10 iterations, respectively. For the soft-impute method, we select $\lambda = 0, K = 4$ and $\lambda = 0, K = 9$ for the 1M and 10M data, respectively. For the restricted Boltzmann machine, we choose 200 hidden units for both the 1M and 10M data.

Compared to the collaborative filtering methods based on matrix factorization/completion, the root mean square error of the two-way logistic model with covariate information is 10.1% smaller than that of the regularized singular value decomposition method, 20.8% smaller than the regression-based latent

factor model, 11.6% smaller than the soft-impute method, and 1.6% smaller than the group-specific recommender system. On the MovieLens 10M data, the root mean square error of the two-way logistic model with covariate information is 7.6% smaller than that of the regularized singular value decomposition method, 5.4% smaller than the regression-based latent factor model, 9.5% smaller than the soft-impute method, and 0.9% smaller than the group-specific recommender system.

Due to the loss of covariate information, the two-way logistic model without covariate information performs 2.2% worse on the MovieLens 1M data and 0.4% worse on the MovieLens 10M data than the same model with covariate information. Compared to the group-specific recommender system (Bi et al. 2017) that also utilizes covariate information to group up the users and the items, the two-way logistic model without covariate information perform 0.5% worse on the MovieLens 1M data, but 0.4% better on the MovieLens 10M data. As for the other matrix-factorization-based collaborative filtering methods, the root mean square error of the two-way logistic model without covariate information is 8.2% smaller than that of the regularized singular value decomposition method, 19.1% smaller than the regression-based latent factor model, and 9.7% smaller than the soft-impute method on the MovieLens 1M data. On the MovieLens 10M data, it is 7.3% smaller than the regularized singular value decomposition method, 5.1% smaller than the regression-based latent factor model, and 9.2% smaller than the soft-impute method. This shows that the two-way logistic model is more consistent with the categorical nature of the ratings.

One the other hand, compared to the restricted Boltzmann machine that also generates categorical predicted ratings, the two-way multinomial logistic model with covariate information also performs 6.3% better on the MovieLens 1M data and 5.8% better on the MovieLens 10M data than the restricted Boltzmann machine. Even without the covariate information, the two-way multinomial logistic model without covariate information also performs 4.3% better on the MovieLens 1M data and 5.4% better on the MovieLens 10M data than the restricted Boltzmann machine.

From the above comparison, we observe that the improvement the two-way logistic model with covariate information over the same model without covariate information is more significant for the MovieLens 1M data than the MovieLens 10M data. This is consistent with the fact that the MovieLens 1M data contain more covariate information than the MovieLens 10M data. In addition, the improvement of both the two-way logistic model with and without covariate information is generally more significant in the MovieLens 1M data than in the MovieLens 10M data. This is because in both the MovieLens 1M and 10M data, each user rates about 150 movies on average, while the average number of ratings for the movies is 250 for the 1M data and 900 for the 10M data. The availability of more past ratings makes predicting the new ratings easier on the MovieLens 10M data than the 1M data.

Table 5. Comparison of the methods on MovieLens 1M and 10M data.

	RSVD	RBM	RB	SI	GS	TML	TMLCI
1M	1.0552	1.0128	1.1974	1.0737	0.9644	0.9692	0.9487
10M	0.9966	0.9772	0.9737	1.0177	0.9295	0.9243	0.9207

6. Conclusion

In this article, we propose a two-way multinomial logistic model for recommender systems on categorical ratings. Specifically, instead of factorizing the ratings immediately into the product of the latent factors of the users and the items, we use a multiclass logistic model where the probability mass function of the ratings of a user on an item is factorized into a function of the latent factor of that user and item. Each user and each item have latent factors representing their a priori absolute preference on giving/receiving a specific rating, and the actual ratings are generated from a probability distribution of the ratings determined by the user's and item's latent factors together with their covariate information through a logistic function.

Compared to the existing collaborative filtering methods, the proposed method can incorporate both the covariate information and the latent factors of the users and the items uniformly in the model. Compared to the matrix factorization/completion based collaborative filtering methods, where the ratings are treated as the discretization of the inner products of the latent factors from the users and the items, we treat the possible ratings as categorical random variables where the probability mass function is determined by the latent factors and the covariate information of the users and the items in the proposed method.

These give the proposed method an advantage in handling categorical ratings and the cold-start problems. As shown by the numerical experiments, the proposed method performs significantly better than the restricted singular value decomposition (Koren, Bell, and Volinsky 2009), the soft-impute matrix completion method (Mazumder, Hastie, and Tibshirani 2010), the regression-based latent factor models (Agarwal and Chen 2009), the restricted Boltzmann machine (Salakhutdinov, Mnih, and Hinton 2007), and the group-specific recommender system (Bi et al. 2017) on the simulated binary rating data under different missing rates and under different data sizes, the simulated multicategorical rating data under different missing rates and in the cold-start problem, and the MovieLens 1M and 10M data.

Supplementary Materials

Additional data on comparison of different methods for different missing rates are available in supplemental material.

Acknowledgments

We greatly appreciate the comments and suggestions provided by the three anonymous reviewers, which improve the manuscript significantly.

Funding

We would like to acknowledge support for this project from the National Science Foundation grants DMS-1415308, DMS-1613190, DMS-1821198, CPS-1329991 and AFOSR grant FA9550-15-1-0059.

References

Acıar, S., Zhang, D., Simoff, S., and Debenham, J. (2006), "Recommender System Based on Consumer Product Reviews," in *Proceedings of the 2006 IEEE/WIC/ACM International Conference on Web Intelligence, WI '06*, Washington, DC, USA, IEEE Computer Society, pp. 719–723. [396]

Adomavicius, G., and Tuzhilin, A. (2005), "Toward the Next Generation of Recommender Systems: A Survey of the State-of-the-Art and Possible Extensions," *IEEE Transactions on Knowledge and Data Engineering*, 17, 734–749. [396]

Agarwal, D., and Chen, B.-C. (2009), "Regression-Based Latent Factor Models," in *Proceedings of the 15th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, KDD '09*, New York, NY, USA, ACM, pp. 19–28. [397,398,399,400,403]

Aggarwal, C. C. (2016), *Recommender Systems*, Cham: Springer. [396]

Agresti, A., and Kateri, M. (2011), "Categorical Data Analysis," in *International Encyclopedia of Statistical Science*, ed. M. Lovric, Berlin, Heidelberg: Springer, pp. 206–208, DOI: 10.1007/978-3-642-04898-2_161. [396]

Bell, R. M., and Koren, Y. (2007), "Lessons From the Netflix Prize Challenge," *SIGKDD Explorations Newsletter*, 9, 75–79. [396]

Bi, X., Qu, A., Wang, J., and Shen, X. (2017), "A Group-Specific Recommender System," *Journal of the American Statistical Association*, 112, 1344–1353. [396,397,398,399,400,402,403]

Blanco-Fernandez, Y., Pazos-Arias, J. J., Gil-Solla, A., Ramos-Cabrer, M., and Lopez-Nores, M. (2008), "Providing Entertainment by Content-Based Filtering and Semantic Reasoning in Intelligent Recommender Systems," *IEEE Transactions on Consumer Electronics*, 54, 727–735. [396]

Bobadilla, J., Ortega, F., Hernando, A., and Alcalá, J. (2011), "Improving Collaborative Filtering Recommender System Results and Performance Using Genetic Algorithms," *Knowledge-Based Systems*, 24, 1310–1316. [396]

Cacheda, F., Carneiro, V., Fernández, D., and Formoso, V. (2011), "Comparison of Collaborative Filtering Algorithms: Limitations of Current Techniques and Proposals for Scalable, High-Performance Recommender Systems," *ACM Transactions on the Web*, 5, 2:1–2:33. [396]

Cron, A., Zhang, L., and Agarwal, D. (2014), "Collaborative Filtering for Massive Multinomial Data," *Journal of Applied Statistics*, 41, 701–715. [397]

Davidson, J., Liebal, B., Liu, J., Nandy, P., Van Vleet, T., Gargi, U., Gupta, S., He, Y., Lambert, M., Livingston, B., and Sampath, D. (2010), "The YouTube Video Recommendation System," in *Proceedings of the Fourth ACM Conference on Recommender Systems, RecSys '10*, New York, NY, USA, pp. 293–296. [396]

Ekstrand, M. D., Riedl, J. T., and Konstan, J. A. (2011), "Collaborative Filtering Recommender Systems," *Foundations and Trends in Human-Computer Interaction*, 4, 81–173. [396]

Ghosh, S., Mundhe, M., Hernandez, K., and Sen, S. (1999), "Voting for Movies: The Anatomy of a Recommender System," in *Proceedings of the Third Annual Conference on Autonomous Agents, AGENTS '99*, New York, NY, pp. 434–435. [396]

Goldberg, K., Roeder, T., Gupta, D., and Perkins, C. (2000), "Eigentaste: A Constant Time Collaborative Filtering Algorithm," *Information Retrieval*, 4, 133–151. [397]

Gomez-Urbe, C. A., and Hunt, N. (2015), "The Netflix Recommender System: Algorithms, Business Value, and Innovation," *ACM Transactions on Management Information Systems*, 6, 13:1–13:19. [396]

Gunawardana, A., and Shani, G. (2015), "Evaluating Recommender Systems," in *Recommender Systems Handbook*, eds. F. Ricci, L. Rokach, and B. Shapira, Boston, MA: Springer, pp. 265–308. [398]

Harper, F. M., and Konstan, J. A. (2015), "The MovieLens Datasets: History and Context," *ACM Transactions on Interactive Intelligent Systems*, 5, 19:1–19:19. [396,397,402]

Iaquinta, L., De Gemmis, M., Lops, P., Semeraro, G., Filannino, M., and Molino, P. (2008), "Introducing Serendipity in a Content-Based Recommender System," in *Eighth International Conference on Hybrid Intelligent Systems, HIS'08*, IEEE, pp. 168–173. [396]

Kim, Y. S., Yum, B.-J., Song, J., and Kim, S. M. (2005), "Development of a Recommender System Based on Navigational and Behavioral Patterns of Customers in e-Commerce Sites," *Expert Systems With Applications*, 28, 381–393. [396]

Koenigstein, N., Dror, G., and Koren, Y. (2011), "Yahoo! Music Recommendations: Modeling Music Ratings With Temporal Dynamics and Item Taxonomy," in *Proceedings of the Fifth ACM Conference on Recommender Systems, RecSys '11*, New York, NY, USA, pp. 165–172. [396]

Koren, Y. (2010), "Factor in the Neighbors: Scalable and Accurate Collaborative Filtering," *ACM Transactions on Knowledge Discovery From Data*, 4, 1:1–1:24. [397]

- Koren, Y., Bell, R., and Volinsky, C. (2009), "Matrix Factorization Techniques for Recommender Systems," *Computer*, 42, 30–37. [396,397,399,400,403]
- Levi, A., Mokryn, O., Diot, C., and Taft, N. (2012), "Finding a Needle in a Haystack of Reviews: Cold Start Context-Based Hotel Recommender System," in *Proceedings of the Sixth ACM Conference on Recommender Systems, RecSys '12*, New York, NY, pp. 115–122. [396]
- Lops, P., De Gemmis, M., and Semeraro, G. (2011), "Content-Based Recommender Systems: State of the Art and Trends," in *Recommender Systems Handbook*, eds. F. Ricci, L. Rokach, B. Shapira, and P. Kantor, Boston, MA: Springer, pp. 73–105. [396]
- Lu, J., Wu, D., Mao, M., Wang, W., and Zhang, G. (2015), "Recommender System Application Developments: A Survey," *Decision Support Systems*, 74, 12–32. [396]
- Mazumder, R., Hastie, T., and Tibshirani, R. (2010), "Spectral Regularization Algorithms for Learning Large Incomplete Matrices," *Journal of Machine Learning Research*, 11, 2287–2322. [396,397,398,399,400,403]
- Melville, P., Mooney, R. J., and Nagarajan, R. (2002), "Content-Boosted Collaborative Filtering for Improved Recommendations," in *Eighteenth National Conference on Artificial Intelligence*, Menlo Park, CA, USA, American Association for Artificial Intelligence, pp. 187–192. [396,397]
- Miller, B. N., Albert, I., Lam, S. K., Konstan, J. A., and Riedl, J. (2003), "MovieLens Unplugged: Experiences With an Occasionally Connected Recommender System," in *Proceedings of the 8th International Conference on Intelligent User Interfaces, IUI '03*, New York, NY, ACM, pp. 263–266. [396,397,402]
- Mirbakhsh, N., and Ling, C. X. (2015), "Improving Top-N Recommendation for Cold-Start Users via Cross-Domain Information," *ACM Transactions on Knowledge Discovery From Data*, 9, 33:1–33:19. [396]
- Mooney, R. J., and Roy, L. (2000), "Content-Based Book Recommending Using Learning for Text Categorization," in *Proceedings of the Fifth ACM Conference on Digital Libraries, DL '00*, New York, NY, pp. 195–204. [396]
- Nguyen, A.-T., Denos, N., and Berrut, C. (2007), "Improving New User Recommendations With Rule-Based Induction on Cold User Data," in *Proceedings of the 2007 ACM Conference on Recommender Systems, RecSys '07*, New York, NY, pp. 121–128. [397]
- Oghina, A., Breuss, M., Tsagkias, M., and De Rijke, M. (2012), "Predicting IMDB Movie Ratings Using Social Media," in *Advances in Information Retrieval*, Lecture Notes in Computer Science, Berlin, Heidelberg: Springer, pp. 503–507. [396,397]
- Park, S.-T., Pennock, D., Madani, O., Good, N., and DeCoste, D. (2006), "Naïve Filterbots for Robust Cold-Start Recommendations," in *Proceedings of the 12th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, KDD '06*, New York, NY, USA, pp. 699–705. [397]
- Resnick, P., and Varian, H. R. (1997), "Recommender Systems," *Communications of the ACM*, 40, 56–58. [396]
- Ricci, F., Rokach, L., and Shapira, B. (2011), "Introduction to Recommender Systems Handbook," in *Recommender Systems Handbook*, Boston, MA: Springer, pp. 1–35. [396]
- Salakhutdinov, R., Mnih, A., and Hinton, G. (2007), "Restricted Boltzmann Machines for Collaborative Filtering," in *Proceedings of the 24th International Conference on Machine Learning, ICML '07*, New York, NY, USA, pp. 791–798. [396,397,398,399,400,403]
- Spall, J. C. (2005), *Introduction to Stochastic Search and Optimization: Estimation, Simulation, and Control* (Vol. 65), Hoboken, NJ: Wiley. [399]
- Srebro, N., Alon, N., and Jaakkola, T. S. (2005), "Generalization Error Bounds for Collaborative Prediction With Low-Rank Matrices," in *Advances in Neural Information Processing Systems* (Vol. 17), eds. L. K. Saul, Y. Weiss, and L. Bottou, Cambridge, MA: MIT Press, pp. 1321–1328. [396]
- Stern, D., Herbrich, R., and Graepel, T. (2009), "Matchbox: Large Scale Bayesian Recommendations," in *Proceedings of the 18th International World Wide Web Conference*. [397]
- Zhu, Y., Shen, X., and Ye, C. (2016), "Personalized Prediction and Sparsity Pursuit in Latent Factor Models," *Journal of the American Statistical Association*, 111, 241–252. [397]