



Contents lists available at ScienceDirect

Physica A

journal homepage: www.elsevier.com/locate/physa

Q1 Network-based recommendation algorithms: A review

Q2 Fei Yu^a, An Zeng^{a,b}, Sébastien Gillard^a, Matúš Medo^{a,*}

^a Department of Physics, University of Fribourg, 1700 Fribourg, Switzerland

^b School of Systems Science, Beijing Normal University, Beijing, PR China

HIGHLIGHTS

- Networks and algorithms based on them are often applied in recommendation.
- We motivate the use of network-based recommendation methods.
- We introduce a comprehensive set of networks-based recommendation methods.
- We use several performance metrics and a robust approach to choose method parameters.
- We compare the methods' performance on three distinct input datasets.

ARTICLE INFO

Article history:

Received 14 September 2015

Received in revised form 11 January 2016

Available online xxxx

Keywords:

Information filtering

Recommender systems

Complex networks

Random walk

ABSTRACT

Recommender systems are a vital tool that helps us to overcome the information overload problem. They are being used by most e-commerce web sites and attract the interest of a broad scientific community. A recommender system uses data on users' past preferences to choose new items that might be appreciated by a given individual user. While many approaches to recommendation exist, the approach based on a network representation of the input data has gained considerable attention in the past. We review here a broad range of network-based recommendation algorithms and for the first time compare their performance on three distinct real datasets. We present recommendation topics that go beyond the mere question of which algorithm to use – such as the possible influence of recommendation on the evolution of systems that use it – and finally discuss open research directions and challenges.

© 2016 Published by Elsevier B.V.

1. Introduction

The rapid development of the Internet has a great impact on our daily lives and has significantly changed the ways in which we obtain information. Movie fans, instead of going to a physical shop to buy or rent a DVD, can now use one of the many online movie-on-demand or rental services to watch the movie they want. Online services have similarly simplified our access to books and music. The same thing happens to our social lives: instead of going to bars to meet with old and possibly also new friends, we now have multiple online social networks which allow us to communicate with friends as well as to find new ones. However, the convenience brought by the Internet comes with the burden to choose from the immense number of possibilities – which movie to watch, which song to hear, whose Tweets to read – which has become to known as the information overload problem [1].

Since it is often impossible for a person to evaluate all the available possibilities, the need has emerged for automated systems that would help to identify the potentially interesting and valuable candidates for any individual user. Many

* Corresponding author.

E-mail address: matus.medo@unifr.ch (M. Medo).

<http://dx.doi.org/10.1016/j.physa.2016.02.021>

0378-4371/© 2016 Published by Elsevier B.V.

information filtering techniques have been proposed to meet this challenge [2]. One representative method is the search engine which returns the most relevant web pages based on the search keywords provided by the users [3]. Though effective and commonly used, search engines have two main drawbacks. First, they require the users to specify the keywords describing the contents that they are interested in, which is often a difficult task, especially when one has little experience with a given topic or, even, when one does not know what they are looking for. Second, search results are not personalized which means that every user providing the same keywords obtains the same results (this problem can be corrected by assessing the individual's history of searches). This is crucial because the tastes and interests of people are extraordinary diverse and ignoring them is likely to lead to inferior filtering performance.

The second class of information filtering techniques, recommender systems, overcomes the above-mentioned problems. The goal of a recommender system is to use data on users' past interests, purchase behavior, and evaluations of the consumed content, to predict further potentially interesting items for any individual user [4]. These data typically take form of ratings given by users to items in an integer rating scale (most often 1–5 stars where more stars means better evaluation) but it can also be of so-called unary kind where a user is connected with an item only if the user has purchased, viewed, or otherwise collected. Since user tastes and interests are included in the input data, recommendations can be obtained without providing any search queries or keywords. The choice of items for a given user builds on the user's past behavior which ensures that the recommendation is personalized. However, the degree of personalization can be harmed by excessive focus on recommendation accuracy [5,6].

Collaborative filtering is perhaps the most usual approach to recommendation [7,8]. User-based collaborative filtering evaluates the similarity of users and recommends items that have been appreciated by users who are similar to a target user for whom the recommendations are being computed (analogously, item-based approach builds on evaluating the similarity of items). Other techniques include content-based analysis [9], spectral analysis [10], latent semantic models [11], matrix factorization [12], and social recommendation [13,14]. The last-mentioned class of algorithms has recently gained popularity because of contributing importantly to the winning solution [15] in the Netflix prize contest [16]. See Refs. [17–22] for a current review of various aspects of the field of recommendation.

While most recommender systems act on data with ratings, unary data without ratings are the basis for a class of physics-inspired recommendation algorithms. These algorithms represent the input data with a bipartite user–item network where users are connected with the items that they have collected (more information on complex networks and their use for analyzing and modeling real systems can be found in Refs. [23–25]; bipartite user–item networks in particular are discussed in Refs. [26,27]). Classical physics processes such as random walk [28] and heat diffusion [29] can be then employed on the network to obtain recommendations for individual users. See Ref. [22] for a review of the basic ideas in network-based recommendation and ranking algorithms. Many variants and improvements of the originally proposed algorithms have been subsequently published and their scope has been extended to, for example, the link prediction problem [30,31] and the prediction of future trends [32].

In this review, we select a comprehensive group of recommendation algorithms that act on unary data and compare them for the first time using several recommendation performance metrics and various datasets that differ in their basic properties such as size and sparsity. After introducing the algorithms and the evaluation procedure in Section 2, we present the results in Section 3. In this section, we focus in particular on evaluating the contribution of additional parameters that are used by most of the algorithms to improve their performance and make it possible to adjust the algorithm to a particular dataset. In Section 4, we discuss several questions that are not directly related to recommendation algorithms. In particular, we expand considerably the findings presented in Refs. [33,34] that can be used to further improve accuracy and diversity of recommendations by limiting the number of users to whom each individual item can be recommended. Finally in Section 5, we summarize the main conclusions of this review and outline the major research directions for the future.

2. Methods

In this section, we describe the notation, benchmark datasets, recommendation methods, and the evaluation procedure and metrics that are used in this review.

2.1. Data and notation

The input data for a recommender system typically consists of past activity records of users. We confine ourselves to the simplest case where the past record for each user is represented by the set of items collected by this user. Further information, such as the time when individual items have been collected or personal information about the user (gender, age, nationality, and so forth) is not required. The input data can be effectively represented by a bipartite user–item network where a user and item node are connected when the corresponding user has collected the given item. In the case when users also rate the collected items, we represent with links only those collected items whose rating is greater than or equal to a chosen threshold rating.

The number of user and item nodes in the network is U and I , respectively. The total number of links in the network is L . In mathematical notation, we speak of the bipartite graph $G(U, I, L)$ where U, I, L is the set of users, items, and links, respectively, and $U := |U|, I := |I|, L := |L|$. To improve the clarity of our notation, we use Latin letters i, j to label user nodes and Greek letters α, β to label item nodes. The degree of user i and item α is labeled as k_i and k_α , respectively. The

Download English Version:

<https://daneshyari.com/en/article/7377991>

Download Persian Version:

<https://daneshyari.com/article/7377991>

[Daneshyari.com](https://daneshyari.com)