

A Survey on Text Based Recommendation System

Asst. Prof. Pandey Nivedita^{*1}, Mr. Swapnil Mankar²

^{*1}Department of CSE VTU University/East Point College Engineering and Technology, Bangalore, Karnataka,
India

²R and D Department Bangalore, Karnataka, India

ABSTRACT

A wide range of textual material is produced by numerous websites on the Internet, including news, academic publications, ebooks, personal blogs, and user reviews. These websites include so much textual content that it can be difficult for users to find the information they need. To resolve this problem, The creation of "Text-based Recommendation Systems (RS)" is underway. They are the systems that can use text as their main feature to quickly find the pertinent information. There are numerous methods available for developing and assessing such systems. Although several surveys list the general characteristics of recommendation systems, a thorough literature assessment on text-based recommendation systems is still lacking. We examine the most recent research on text-based RS in this publication. The four primary facets of text-based recommendation systems employed in the reviewed literature are primarily covered by this survey. Datasets, feature extraction methods, computational frameworks, and evaluation measures make up the aspects. Publicly accessible datasets are thoroughly evaluated in this work since benchmark datasets are essential to any research. But in order to make these publicly available and proprietary datasets' properties more known to new researchers, we have combined them together. Additionally, the use of the feature extraction techniques from the text in the development of text-based RS is covered. The usage of these properties in various computational methods is later described. Some evaluation metrics are used to assess these systems. According to the report, Word Embedding is the most popular feature selection method currently being employed in research. We can also conclude that combining text features with other features improves the accuracy of the recommendations. The study emphasizes the fact that the majority of the effort focuses on English textual data and that the most common area is news recommendation.

Keywords—text based recommendation system. Filtering

I. INTRODUCTION

Recommender Systems (RSs) gather data on consumer preferences for a range of goods, including films, music, books, jokes, gadgets, software, websites, vacation spots, and e-learning materials. The data may be gathered directly (commonly through user ratings), or inadvertently (usually by surveillance of user behavior, such as music listened to, programmes downloaded, websites viewed, and books read). User demographics (such as gender, age, and nationality) may be used by RS. Web 2.0 regularly makes use of social media statistics

including follows, followed, twits, and postings. More and more people are using data from the Internet of Things, such as GPS coordinates, RFID, and live health signals.

RS employs a range of information sources to create forecasts and product recommendations for customers. When giving proposals, they try to strike a balance between elements like stability, dispersity, innovation, and correctness. Collaborative filtering (CF) techniques are important in the recommendation process, even if they are commonly used in conjunction with other filtering strategies like content-based, knowledge-based, or social ones. CF is founded on how people have made decisions throughout history: in addition to our own experiences, we also rely our choices on information and understanding that comes to each of us from a sizable network of acquaintances.

The collaborative, content-based, and demographic filtering types that were most popular at the start of the RS were discussed. Breese et al. assessed the CF predictive accuracy of several algorithms; thereafter, the classical publication defines the framework for Collaborative Filtering RS assessment. The importance of RS hybrid techniques, which combine many techniques to gain the benefits of each, has been demonstrated by the progress of the discipline. The hybrid RS has been the subject of a survey. The use of social-filtering, a method that has gained popularity recently through social networks, is not addressed, though.

At the start of the RS, the neighborhood-based CF has been the most widely used recommendation approach; Herlocker et al. present a set of criteria for creating neighborhood-based prediction systems. In their overview of the RS field, Adomavicius and Tuzhilin [3] highlight the most challenging areas where RS researchers should concentrate on the "next generation of RS": limited content analysis and overspecialization in content-based methods, cold-start and sparsity in CF methods, model-based techniques, non intrusiveness, flexibility (real-time customization), etc.

Our survey seeks to advance the evolution of the RS, moving from a first phase based on the traditional Web to the present second phase based on social Web, which is currently moving to a third phase (Internet of things), as opposed to the existing surveys, which concentrate on the most pertinent methods and algorithms of the RS field. In order to be helpful to the new readers of the RS area, we have incorporated certain conventional themes within this survey, including RS foundations, the k-Nearest Neighbours method, cold-start difficulties, similarity metrics, and RS evaluation. The remainder of the paper discusses brand-new subjects that are not included in current surveys. Advanced readers in RS will learn about social information (social filtering: followers, followed, trust, reputation, credibility, content-based filtering of social data; social tagging and taxonomies) in depth through this survey, as well as how to explain recommendations to groups of users. This survey will be helpful to readers who are interested in brand-new and upcoming applications because it provides information on the most recent works in location-aware RS trends and bio-inspired techniques. They will also learn about certain crucial topics, like teleoperation, telepresence, privacy, security, P2P information, and the utilization of the Internet of Things (RFID data, health parameters, surveillance data, etc.).

We provide a clear explanation of the process used to choose the most important papers in the RS field. The methods, algorithms, and models for recommending material based on information from the traditional web, including ratings, demographic information, and item data (CF, demographic filtering, content-based filtering, and hybrid filtering) are described. Measures for assessing the accuracy of the RS predictions and recommendations are described. The utilization of social information from Web 2.0 for recommendations through ideas like trust, reputation, and credibility is demonstrated. We will also discuss methods for obtaining social information based on content (such as tags and postings).

Online data has increased dramatically as a result of digital technological developments, particularly with the launch of smartphones. Social media platforms like Twitter and Facebook are important sources of data generation. Additionally, websites that answer questions, like Quora and Stack overflow, are rapidly adding data to this pool. Furthermore, during the past few years, both the trend of personal blogging and the quantity of publications have grown significantly. Users' lives have been impacted by the introduction of digitalization in both positive and negative ways. The data is readily and immediately accessible, which is a good thing. Recommendation systems have been created to address the issue of locating the most pertinent data among the overabundance of data [1].

II. SURVEY

It is a special collection of tools and methods that recommend to a user certain things that might be of interest to them. A recommendation system keeps track of a customer's profile and suggests a good or service based on their interests [2]. These recommendations can be found in any domain, from a news story to read to a web service to employ in software. The recommendation problem is divided into two parts, namely (i) calculating the value of prediction for each item and (ii) ranking these things according to their prediction value. And there are numerous methods to complete this work. They are most well-liked by (CF) and (CB) [1]. In CF, the user is presented with a new item via the recommendation system, which the user then consumes. In contrast, the recommendation system for CB suggests a new item based on its content and characteristics. In other words, the recommended item will share characteristics with the goods that the same person has already consumed. But recently, both of these methods have been combined and have produced encouraging outcomes. As any sort of data, including photos, text, numbers, and others, can be used, any type of data could be put into the recommendation system. The textual data makes up a sizeable share of the available data kinds. For instance, the main sources of this work include news, research publications, blogs, and various reports. The amount of textual data presents similar difficulties to those previously described for data in general. As a result, academics have also focused on textual data recommendation algorithms to discover the most pertinent content. Textual recommendations can be found in a variety of contexts, including news, articles, blogs, books, and movies. These text-based domains have been the focus of research for several decades [6]. There are unique problems in every domain [7]. As a result, each domain uses somewhat distinct methodologies. An overview of all such effort is therefore absolutely necessary. In practically every area of computer science, deep learning models have recently taken the role of conventional algorithms. Similarly, this new tendency has been embraced by the field of recommendation systems. For textual recommendations, there are several different neural network topologies that can be employed.

This assessment also provides a summary of the most recent methods used in textual recommendation systems. This page also compiles the evaluation methods for textual recommendation systems. Traditional assessment measures from the discipline of information retrieval are primarily used for RS. However, textual RS also makes use of some unusual measures like specificity and diversity [12]. They are most well-liked by (CF) and (CB) [1]. In CF, a new item that is consumed by users who are similar to the user is recommended to them. In contrast, the recommendation system for CB suggests a new item based on its content and characteristics. In other words, the suggested item will share characteristics with the foods you've already eaten.

But recently, both of these methods have been combined and have produced encouraging outcomes. As any sort of data, including photos, text, numbers, and others, can be used, any type of data could be put into the recommendation system. The textual data makes up a sizeable share of the available data kinds. For instance, the main sources of this work include news, research publications, blogs, and various reports. The problems presented by the abundance of textual data are similar to those raised earlier for data in general. To locate the most pertinent material, academics have also focused on textual data recommendation algorithms. News recommendations, article/blog recommendations, book/movie recommendations, etc. are a few examples of textual recommendation domains. Work is being done by researchers on these text-based

III. OBJECTIVES FOR RESEARCH

The primary goals of this study are listed below.

- One of the main goals of the study is to provide the groundwork for future research on text-based recommendation systems, which has grown to be a significant research paradigm. The accomplishments related to textual data recommendations will be highlighted in the essay.
- To construct assessment metrics used in the evaluation of text-based recommendation systems;
- To gain understanding of how to extract meaningful features using various methodologies.
- To provide a general overview of the attributes present in the datasets utilized in textual RS. The purpose of this study is to provide a head start for researchers who wish to continue their work in the area of textual recommendation systems.

IV. EXISTING SYSTEM

There are numerous surveys that explore the various facets of recommendation systems. Some of them are all-encompassing, while others are focused on a particular topic, like in Chen et al. which lists all recommendation systems that rely on user reviews. A brief history of content-based recommendation systems was offered in another study that analyzed recent trends in content-based recommendation systems. The most recent trends were discussed in terms of data and algorithms. The poll said that Link Open Data (LOD), which is used to obtain additional meta-data attributes for an item, is now commonly used for data purposes.

In addition, more content is acquired through forums, user reviews, and tags; this content is classed as User Generated Data (UGD), and heterogeneous information networks are also explored as a source of content. The following methods were emphasized in terms of algorithms; Deep learning approaches, novel metadata encoding, word and document embeddings to find latent features in the text, and meta-path based methods are some of the methods used to represent a path that corresponds to a relationship between two things.

A different study concentrated exclusively on those RS that used ontologies in the creation of e-learning RS. All of the potential types of RS were briefly covered at the outset of the study. The writer summarizes all the papers at hand in terms of the ontologies employed, the ontology representation language, and suggests learning resources after giving a thorough explanation of ontologies and e-learning systems. A thorough analysis of the recommendation system that is solely based on user reviews (UR) or whose effectiveness has been enhanced by UR was conducted in another survey . The broad introduction to RS and its fundamental methodologies (Content-based, Rating-based Collaborative Filtering, and Preference-based Product Ranking)

were provided in the first section of the survey.

The second section of the poll covered the components of the reviews utilized in RS, including frequent terms, review themes, feature opinions, contextual opinions, comparative opinions, review emotions, and review helpfulness. All of the studies that used UR for user and product profiling for recommendations were thoroughly discussed in the following two sections of the survey, and in the final section, they discussed the practical implications of their findings for five dimensions: data quality, addition of new users, advancement of algorithms, profile-building, and product domain. One such recent study provides information on the use of text mining techniques in various recommendation systems.

Although the paper's primary focus is deep learning approaches and text-based RS is not specifically mentioned, it does cover all work on textual data that has been done using the DL methodology. A paper published by Batmaz et al. compiles deep learning methods utilized in RS. It gives a quick overview of the development methodologies but focuses more on the problems and difficulties that DL can help solve. Additionally, a structured form of the domains in which these models are used is also provided. There are also some studies that focus on particular domains, such as the one reported by [7].

Despite the fact that the text-based approaches discussed in this are quite brief and simple, it nonetheless provides an overview of the trend among those who work on News RS. The paper provides a summary of News RS algorithmic approaches, difficulties, and evaluation criteria. This work also contains condensed summaries of popular and openly accessible data sets. Another Study covers the difficulties and techniques in the news industry, although it is more general and seldom ever specific about any design on pure textual data.

V. METHODOLOGY

There are numerous surveys that explore the various facets of recommendation systems. Some of them are all-encompassing, while others are focused on a particular topic, like in Chen et al. , which lists all recommendation systems that rely on user reviews. A brief history of content-based recommendation systems was offered in another study that analyzed recent trends in content-based recommendation systems. The most recent trends were discussed in terms of data and algorithms. The poll said that Link Open Data (LOD), which is used to obtain additional meta-data attributes for an item, is now commonly used for data purposes.

All of the potential types of RS were briefly covered at the outset of the study. The writer summarizes all the papers at hand in terms of the ontologies employed, the ontology representation language, and suggested learning resources after giving a thorough explanation of ontologies and e-learning systems. A thorough analysis of the recommendation system that is solely based on user reviews (UR) or whose effectiveness has been enhanced by UR was conducted in another survey . The broad introduction to RS and its fundamental methodologies (Content-based, Rating-based Collaborative Filtering, and Preference-based Product Ranking) were provided in the first section of the survey. The second section of the poll covered the components of the reviews utilized in the study, including frequent terms, review themes, feature opinions, contextual opinions, comparative opinions, review emotions, and review helpfulness.

All of the studies that used UR for user and product profiling for recommendations were thoroughly discussed in the following two sections of the survey, and in the final section, they discussed the practical implications of their findings for five dimensions: data quality, addition of new users, advancement of algorithms, profile-building, and product domain. One such recent study provides information on the use of text mining

techniques in various recommendation systems. Deep learning has been more popular recently, even in the field of RS, and virtually everyone is using it to construct RS because of its successful and promising outcomes. All of these works were gathered in an overview on deep learning methods used in RS.

Although the paper's primary focus is deep learning approaches and text-based RS is not specifically mentioned, it does cover all work on textual data that has been done using the DL methodology. A paper published by Batmaz et al. compiles deep learning methods utilized in RS. It gives a quick overview of the development methodologies but focuses more on the problems and difficulties that DL can help solve. Additionally, a structured form of the domains in which these models are used is also provided.

VI. ALGORITHM

THE K NEAREST NEIGHBORS RECOMMENDATION ALGORITHM

The reference algorithm for the collaborative filtering recommendation process is the k Nearest Neighbours (kNN) algorithm. Its main benefits are simplicity and reassuringly reliable results; nonetheless, its biggest drawbacks are limited scaling and susceptibility to sparsity in RS databases. An overview of this algorithm function is given in this section.

The kNN-based CF is conceptually simple, has an easy implementation, and typically yields forecasts and suggestions of high caliber. Due to the high amount of sparsity in RS databases, however, similarity measures frequently run into processing issues (usually caused by a lack of mutual ratings when comparing users and items) and cold start scenarios (users and items with few rankings).

The kNN algorithm's poor scalability is yet another significant issue. The procedure of creating a neighborhood for an active user gets excessively slow as databases (like Netflix) grow larger; the similarity measure needs to be evaluated every time a new user registers in the database. The scalability issue is greatly reduced by the item to item variant of the kNN method [200]. In order to achieve this, neighbors are determined for each item; their top n similarity values are then recorded, and over time, predictions and suggestions are made using the stored data. Although ratings from earlier processing or storage are not included in the recorded information, obsolete information for products is less sensitive than for users.

The creation of measurements to precisely and accurately determine the present similarity for the users (or things) is a recurring theme in CF research. A number of statistical measures, including the Pearson correlation, cosine, constraint Pearson correlation, and mean squared differences, have historically been applied. Metrics have recently been developed to meet the limitations and quirks of RS. The relevance notion was created to give greater weight to users and products that were more pertinent. Additionally, a set of measures was created expressly to work well under cold-start conditions. Based on similarity metrics, the kNN algorithm operates. More information on the existing RS similarity measurements is provided in the following sections. According to each of these factors, the similarity approaches often calculate the similarity between two users, x and y (user to user) users' item ratings. The item to item kNN version computes the similarity between two items i and j.

A formal approach of the kNN algorithm may be found in . In this section, we will provide an illustrative example of this algorithm. The method for making recommendations is based on the following three steps:

- (a) Using the selected similarity measure, we produce the set of k neighbors for the active user a. The k neighbors for a are the nearest k (similar) users to u.

- (b) Once the set of k users (neighbors) similar to active a has been calculated, in order to obtain the prediction of item i on user a , one of the following aggregation approaches is often used: the average, the weighted sum and the adjusted weighted aggregation (deviation-from-mean).
- (c) To obtain the top- n recommendations, we choose the n items, which provide most satisfaction to the active user according to our predictions.

RS have included social information (such as posts, blogs, tags, friends lists, followed and unfollowed individuals), which has become more prevalent as web 2.0 has grown. The RS is enhanced by the added contextual information. Because social information supports traditional memory-based information (user ratings), such as users connected by a network of trust exhibit significantly higher similarity on items and meta-data than non-connected users, the sparsity problem inherent in memory-based RS is improved. Researchers use social information for three main reasons: to produce or suggest new RS, to enhance the accuracy of forecasts and recommendations, and to clarify the most important connections between social information and collaborative processes.

A major subject of research in RS is trust and reputation. The social data now present in RS is directly tied to this topic. The following are the most popular methods for creating trust and reputation measurements: User trust (a): determining a user's trustworthiness based on the other users' explicit information or determining a user's trustworthiness based on implicit data gleaned from a social network. Users can introduce labels linked to objects in the social RS field. Folksonomies are collections of information spaces made up of the triples "user, item, tag." Folksonomies are fundamentally used in two ways:

- (1) to build tag recommendation systems (RS based purely on tags);
- (2) To employ tags to enhance recommendation processes.

Social Filtering

Through the identification of a community network or affinity network using the unique data that users provide (such as messages and web logs), social information can be obtained overtly or implicitly. It is possible to enhance the RS outcomes by employing solely user ratings, thereby fostering an implicit social network. To provide suggestions, one can blend implicit and explicit knowledge sources.

A trust-based CF can make advantage of the explicit social information to enhance the quality of suggestions. Different methods, such as trust propagation mechanisms, a "follow the leader" strategy, personality-based similarity measurements, trust networks distrust analysis, and dynamic trust based on the ant colonies metaphor can be used to produce or utilize trust information.

Through the identification of a community network or affinity network using the unique data that users provide (such as messages and web logs), social information can be obtained overtly or implicitly. It is possible to enhance the RS outcomes by employing solely user ratings, thereby fostering an implicit social network.

To provide suggestions, one can blend implicit and explicit knowledge sources.

A trust-based CF can make advantage of the explicit social information to enhance the quality of suggestions. Different methods, such as trust propagation mechanisms, a "follow the leader" strategy, personality-based similarity measurements, trust networks distrust analysis, and dynamic trust based on the ant colonies metaphor, can be used to produce or utilize trust information.

They therefore propose that using the user's resemblance and familiarity with the people who have rated the things can help with judgment and decision-making. By combining social network data into CF, Fengkun and

Hong created a method to improve the effectiveness of recommendations. They obtained information about users' social network connections and preference ratings from a social networking website. After that, they assessed how well CF performed with various neighbor groups, which combined friends and the closest neighbors. According to Carmagnola et al. , joining a network exposes people to social dynamics that can affect their attitudes, behaviors, and preferences: SoNARS, a recommendation algorithm for content in social RS, is presented. SoNARS seeks out users who are active on social networks.

A recommendation system is what?

Numerous applications exist where websites gather user data and use that data to forecast the preferences of their users. They can then recommend the content they enjoy. Recommender systems are a technique of recommending things and concepts that are comparable to a user's particular way of thinking.

VII. CONTENT-BASED FILTERING

The goal of content-based filtering (CBF) is to suggest to the active user goods that have previously received favorable ratings. It is predicated on the idea that objects with comparable features will receive comparable ratings and .For instance, if a user enjoys a website that has the phrases "car," "engine," and "gasoline," the CBF will suggest further websites on the automotive industry.

As RS incorporates data on objects from users operating in web 2.0 contexts, such as tags, posts, opinions, and multimedia content, CBF is becoming increasingly significant.Overspecialization and insufficient content analysis are two difficult issues for content-based filtering [3]. The first issue is that it is challenging to automatically extract trustworthy data from a variety of resources (such as photographs, video, audio, and text), which can significantly lower the quality of suggestions. The second issue, known as overspecialization, is when users only receive recommendations for products that are extremely similar to the products they liked or preferred. As a result, users miss out on suggestions for products they might like but have never heard of (for example, when a user only receives recommendations for fiction films). The uniqueness of recommendations can be assessed .

The qualities of the products you want to recommend must be retrieved for CBF to function . Depending on the item's domain, a set of attributes is often explicitly established for each one. Classic information retrieval techniques must be employed to automatically construct such attributes in some situations, such as when it is wanted to recommend textual material (e.g., term frequency, inverse document frequency, and normalization to page length).

The CBF mechanism is depicted in Fig. and consists of the following steps: To make recommendations, one must first extract the properties of the products, then compare those features to the preferences of the active user, and finally suggest goods having those characteristics.

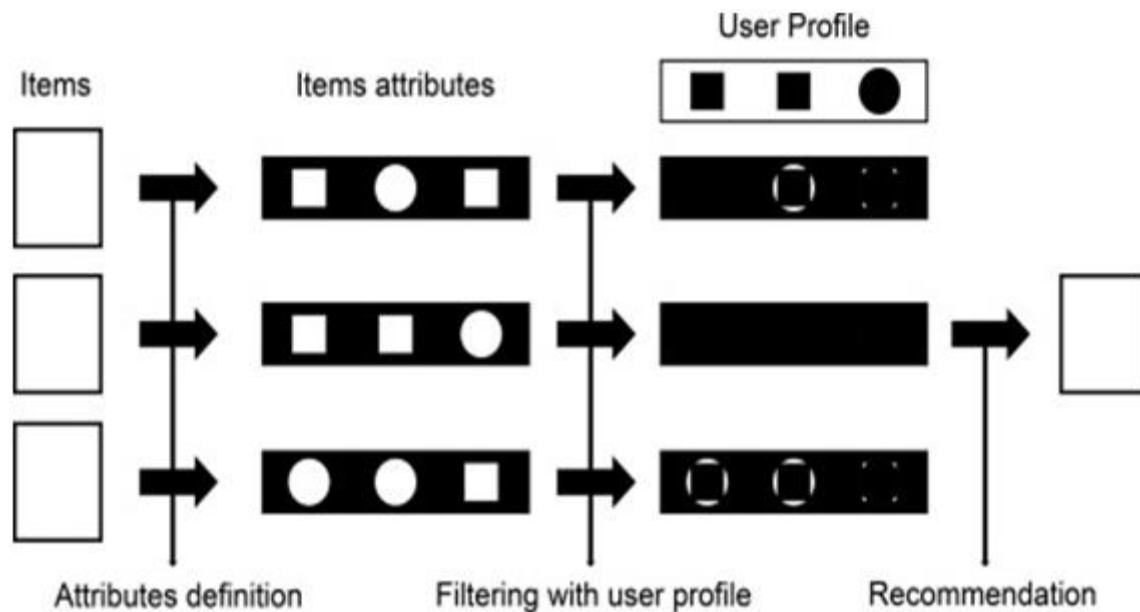


Figure no:1

The main goal of CBF is to ascertain if a user will like a particular thing once the attributes of the objects and user profiles are understood. Traditionally, heuristic techniques or classification algorithms, such as rule induction or nearest neighbors methods, are used to handle this problem. linear classifiers and probabilistic techniques are used in Rocchio's approach. The pure CBF has a number of drawbacks:

- Generating the attributes for things in some domains is a challenging task.
- Because CBF tends to promote the same kinds of products by default, it has an overspecialization problem.
- It is more challenging to get user feedback because, unlike CF, users of CBF don't frequently review the things, making it impossible to know whether the recommendation is accurate.

VIII. COLLABORATIVE BASED FILTERING

In Collaborative Filtering, we tend to find similar users and recommend what similar users like. In this type of recommendation system, we don't use the features of the item to recommend it, rather we classify the users into clusters of similar types and recommend each user according to the preference of its cluster.

There are basically four types of algorithms or say techniques to build Collaborative filtering based recommender systems:

- Memory-Based
- Model-Based
- Hybrid
- Deep Learning

IX. CONCLUSION

Text-based recommendation systems have become more prevalent in the last decade because the internet is generating the bulk of textual data every day over different websites. This study aims to explore text-based

recommendation literature and summarize critical approaches to provide a single platform for the understanding of new researchers. The survey covers four main aspects of a text-based recommendation system. First, what are the fundamental techniques of feature extraction used in text-based recommendation systems. Second, proprietary and publicly available datasets and their details. Third, how such systems are evaluated, what are the most frequently used evaluation metrics, and lastly what algorithmic approaches are opted to formulate the problem. Second, proprietary and publicly available datasets and their details. Third, how such systems are evaluated, what are the most frequently used evaluation metrics, and lastly what algorithmic approaches are opted to formulate the problem.

X. REFERENCES

- [1]. J. Bobadilla, F. Ortega, A. Hernando, and A. Gutiérrez, "Recommender systems survey," *Knowl.-Based Syst.*, vol. 46, pp. 109–132, Jul. 2013.
- [2]. Y. Wang, N. Stash, L. Aroyo, L. Hollink, and G. Schreiber, "Semantic relations for content-based recommendations," in *Proc. 5th Int. Conf. Knowl. Capture (K-CAP)*, 2009, pp. 209–210.
- [3]. Z. Cao, X. Qiao, S. Jiang, and X. Zhang, "An efficient knowledge-graphbased Web service recommendation algorithm," *Symmetry*, vol. 11, no. 3, p. 392, Mar. 2019.
- [4]. X. Kong, M. Mao, W. Wang, J. Liu, and B. Xu, "VOPRec: Vector representation learning of papers with text information and structural identity for recommendation," *IEEE Trans. Emerg. Topics Comput.*, early access, Apr. 26, 2019, doi: 10.1109/TETC.2018.2830698.
- [5]. V. Setty and K. Hose, "Event2Vec: Neural embeddings for news events," in *Proc. 41st Int. ACM SIGIR Conf. Res. Develop. Inf. Retr.*, New York, NY, USA, 2018, pp. 1013–1016.
- [6]. J. Beel, S. Langer, B. Gipp, and A. Nürnberger, "The architecture and datasets of docear's research paper recommender system," *D-Lib Mag.*, vol. 20, no. 11/12, 2014.
- [6]. J. Beel, B. Gipp, S. Langer, and C. Breiteringer, "Research-paper recommender systems: A literature survey," *Int. J. Digit. Libraries*, vol. 17, no. 4, pp. 305–338, Nov. 2016.
- [7]. C. Musto, T. Franza, G. Semeraro, M. de Gemmis, and P. Lops, "Deep content-based recommender systems exploiting recurrent neural networks and linked open data," in *Proc. Adjunct Publication 26th Conf. User Modeling, Adaptation Personalization*, New York, NY, USA, 2018, pp. 239–244.
- [8]. P. Vincent, H. Larochelle, I. Lajoie, Y. Bengio, and P.-A. Manzagol, "Stacked denoising autoencoders: Learning useful representations in a deep network with a local denoising criterion," *J. Mach. Learn. Res.*, vol. 11, no. 12, pp. 3371–3408, 2010.
- [9]. G. de Souza Pereira Moreira, F. Ferreira, and A. M. da Cunha, "News session-based recommendations using deep neural networks," in *Proc. 3rd Workshop Deep Learn. Rec. Syst.*, New York, NY, USA, 2018, pp. 15–23.
- [10]. G. Shani and A. Gunawardana, *Evaluating Recommendation Systems*. Boston, MA, USA: Springer, 2011, pp. 257–297.
- [12]. D. Jannach, L. Lerche, I. Kamehkhosh, and M. Jugovac, "What recommenders recommend: An analysis of recommendation biases and possible countermeasures," *User Model. User-Adapted Interact.*, vol. 25, no. 5, pp. 427–491, Dec. 2015.
- [11]. L. Chen, G. Chen, and F. Wang, "Recommender systems based on user reviews: The state of the art," *User Model. User-Adapted Interact.*, vol. 25, no. 2, pp. 99–154, Jun. 2015.