

Using Bibliometrics and Fuzzy Linguistic Modeling to Deal with Cold Start in Recommender Systems for Digital Libraries

Alvaro Tejeda-Lorente¹(✉), Juan Bernabé-Moreno¹, Carlos Porcel²,
and Enrique Herrera-Viedma¹

¹ Department of Computer Science and Artificial Intelligence,
University of Granada, Granada, Spain
atejeda@decsai.ugr.es

² Department of Computer Science, University of Jaén, Jaén, Spain

Abstract. Every recommender system approach suffers the cold start problem to a greater or lesser extent. To soften this impact, the more common solution is to find the way of populating users profiles either using hybrid approach or finding external data sources. In this paper, we present a fuzzy linguistic approach that using bibliometrics aids to soft or remove the necessity of interaction of users providing them with personalized profiles built beforehand, thus reducing the cold start problem. To prove the effectiveness of the system, we conduct a test involving some researchers, aiming to build their profiles automatically. The results obtained proved to be satisfactory for the researchers.

Keywords: Recommender system · Cold start · Fuzzy linguistic modeling · Digital library

1 Introduction

In the era of Big Data, the amount of information generated on every field in the Web is growing constantly leading to the well known information overload problem [6]. Recommender Systems (RSs) appear as a natural solution to this problem providing personalized recommendations to their users filtering out the non valuable information for them [9]. However, a key feature which defines a recommender system are the user profiles that allow them to provide users with recommendations that are suiting them better. One of the main problems RSs have is the *cold start* problem, that is, when a user or item is new to the system and nothing is known about him/her or it. This is a known problem that has been addressed several times in the literature [23]. However, the importance of providing a suitable solution to this problem is scaling due to the growing presence in the Web of systems requiring personalization as well as real time interaction.

In the academic world, research is a field of main importance. It is an over specialized field where each specialty is quite specific. In a previous work, we presented REFORE, a quality-based fuzzy linguistic recommender system for researchers [24] which main purpose was to aid researchers by keeping them up-to-date regarding the new articles that might be relevant for them. In order to solve the cold start problem, researchers must select the top 5 articles of their profiles that show better the current research topics they are interested on. Afterwards, they should provide them with a linguistic assessment of the importance of each of them. Same process is required for the keywords that define better his current research interest.

To deal with the cold start problem in most of the systems, it requires a previous set up of the user profile done by users themselves, requiring time and a fix and complete profile into the system [17]. REFORE was conceived as a system which delivered to researchers every month an email with the most relevant papers for them from this same month, working as a service to keep the authors up-to-date on their research topics. However, the necessity of a previous and good set up stopped new users to try the tool. Therefore, to reduce the cold start problem removing the interaction of the user to the minimum is basic to provide the system with a better acceptance degree among users. It will allow us to expand REFORE to other possibilities such like being used without any necessity of registration.

The proposed approach provides automatically profiles for users based on their names by extracting the information from their historic research trajectory. This profile is used by REFORE reducing the cold start problem impact that forced users to provide all of this information beforehand.

The aim of this paper is to present a semi automatic fuzzy linguistic solution for the cold start problem in REFORE. This solution that could be applied in whatever bibliographic database from a University Digital Library allowing us to open REFORE to a wider research community by reducing the necessity of a previous profile set up done the users. The major innovations and contributions of the solution include:

1. The provision of automatic profiles extracted from the authors name who are using the system dealing with (incomplete information) and transforming novelty, frequency and quality on users interest.
2. The ability of using the system within less than 1 min without the necessity of registering on it, providing the researchers the possibility of obtaining the more relevant publications for themselves.
3. The system uses fuzzy linguistic modeling to improve the user-system interactions.

The paper is organized as follows: In Sect. 2 the background is presented, that is the basis of recommender systems, fuzzy linguistic modeling and other approach to the cold start problem; Sect. 3 presents the new method for the user profile generation; Sect. 4 addresses the validation of the system, and Sect. 5 offers conclusions based on the study findings.

2 Background

With this section we provide the needed background information to describe our system. First, we will present a short description of recommender systems, then a brief explanation of the cold start problem followed by a description of the fuzzy linguistic modeling.

2.1 Recommender Systems

RSs produce personalized recommendations as output or guide users in a personalized way through a wide range of possible options [2]. Well known examples of successful use of RSs are given in e-commerce [3, 19], health [7] or learning [15]. In order to do that, the system must have knowledge from users. This knowledge can be obtained from different sources and has to be related directly or indirectly with the recommendations provided by the system. That can be done in an *implicit way* through the normal functioning of the system, i.e.: Ratings from a movie, geographical proximity to a shop, preferences regarding tastes, etc., or in a *explicit way* when users are required to provide the information manually to the system [8]. Some systems, as for example some movie recommenders, force users to fill some questionnaires or to rate certain selected movies before any recommendation could be received in order to avoid the cold start problem.

Different categorizations have been proposed for recommender systems based on the approach followed to generate recommendations, being the one who split them on two categories the more extended: content-based and collaborative [9]. Content-based recommender systems are based on the similarity of an user profile with an item profile, meanwhile in collaborative recommender systems the recommendations are generated based on the ratings or behavior provided by other similar users.

On the one hand, collaborative systems [10] use to perform better in some domains adding diversity to the recommendations. However, those systems required an important amount of information gathered from the users behavior making them relatively weaker when dealing with the cold start problem. On the other hand, content-based approaches [20] perform better with new users where their taste is rapidly defined. They have the problem of lacking diversity and serendipity. Each approach has advantages and disadvantages, the combination of the both in a hybrid system tend to mitigate the problems they have [2], e.g.: content-based deal in a better way with the cold start problem, so combining this approach with the collaborative benefits from the advantages of the both. However these recommender systems tend to fail when little is known about users information needs.

2.2 Cold Start Problem

Cold start problem is present in certain information systems where the lack of knowledge affect the system purpose [23]. It is particularly present in recommender systems where the basic functioning is based on the amount of knowledge

accumulated over users or items [17]. We found two main variants: new user cold start problem and new item cold start problem. Since most of the time recommender systems use historical ratings as part of user profiles or item profiles, the issue is present on the both sides.

Different approaches have been trying to deal with the cold start problem with different results. On the one hand, hybrid recommendation approaches, where the mix of different recommendation techniques are used in order to take advantage of each other's strengths [22]. E.g.: softening the problem of collaborative filtering with content-based support [4, 8]. On the other hand, since each new item or user introduced in the system presents a problem, implicitly or explicitly populating profiles is a common solution. User profiles are enriched through information either inferred through some technique or provided by users themselves, meanwhile item profiles are done through rich metadata descriptions or accelerating the rating acquisitions [25]. E.g.: In [17], authors used the binary classifier C4.5 [16] and Naive Bayes algorithm [21] in a previous phase to build user profiles.

Le Hoan Son in [23] present a review of different algorithms and their effectiveness against cold start. A classification into three groups is proposed: *making use of additional data sources*, *selecting the most prominent groups of analogous users* and *enhancing the prediction using hybrid methods*. Results showed a better performing of the algorithm denominated *new heuristic similarity model* [18] which belong to the second group and has no need of additional information.

The basic idea of the most of the solutions is to extract some information not provided in the moment the user profile is set based on the rest of the available information. e.g.: Clustering of users based on common characteristics like geographical information, or adding information from an external source.

2.3 Fuzzy Linguistic Modeling

Information is not always able to be evaluated in a quantitative manner, in some occasions it has to be assessed in a qualitative way. The fuzzy linguistic modeling is based on the concept of *linguistic variable* [26] which has proven good results for modeling qualitative information in many problems [14]. Some classic solutions when it comes to fuzzy linguistic modeling are: classic fuzzy linguistic modeling [1, 26] and ordinal fuzzy linguistic modeling [5].

A typical problem when it comes to fuzzy linguistic modeling is the loss of information that use to happen with approaches like classical and ordinal [26]. In [12] authors present the 2-tuple approach for fuzzy linguistic modeling. It consists on a continuous model of representation of information that allows to reduce the typical information loss.

Let $S = \{s_0, \dots, s_g\}$ be a linguistic term set with odd cardinality, where the mid term represents an indifference value and the rest of the terms are symmetrically related to it. We assume that the semantics of labels are given by means of triangular membership functions and consider all terms distributed on a scale on which a total order is defined [11]. If a symbolic method aggregating linguistic information obtains a value $\beta \in [0, g]$, and $\beta \notin \{0, \dots, g\}$, then an

approximation function is used to express the result in S . β is represented by means of 2-tuples (s_i, α_i) , where $s_i \in S$ represents the linguistic label of the information, and α_i is a numerical value expressing the value of the translation from the original result β to the closest index label, i , in the linguistic term set ($s_i \in S$). This model defines $\Delta(\beta) = (s_i, \alpha)$ and $\Delta^{-1}(s_i, \alpha) = \beta \in [0, g]$ as a set of transformation functions between numeric values and 2-tuples.

In order to establish the computational model a definition of a negation, comparison and aggregation operators is needed. Using the transformation functions above described Δ and Δ^{-1} that avoid the loss of information, any of the existing aggregation operators can be easily extended for dealing with linguistic 2-tuples [12].

When modeling the information a problem arises if different uncertainty degrees on the phenomenon are perceived. In order to deal with that matter an important parameter to determine known as the “granularity of uncertainty” is needed, i.e., the cardinality of the linguistic term set S [13]. In [13] a multi-

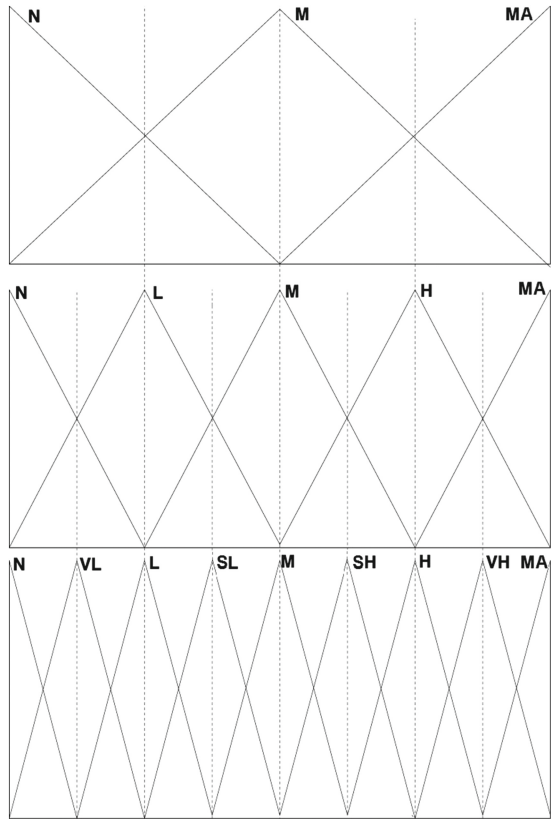


Fig. 1. Linguistic hierarchy of 3, 5 and 9 labels.

granular 2-tuple fuzzy linguistic modeling based on the concept of linguistic hierarchy is proposed.

A *Linguistic Hierarchy*, LH , is a set of levels $l(t, n(t))$, where each level t is a linguistic term set with different granularity $n(t)$ from the remaining hierarchy levels. The levels are ordered according to their granularity, i.e., a level $t + 1$ provides a linguistic refinement of the previous level t . We can define a level from its predecessor level as: $l(t, n(t)) \rightarrow l(t + 1, 2 \cdot n(t) - 1)$. A graphical example of a three level linguistic hierarchy is shown in Fig. 1. Using this LH , the linguistic terms in each level are the following:

- $S^3 = \{a_0 = \text{Null} = N, a_1 = \text{Medium} = M, a_2 = \text{Maximum} = MA\}$.
- $S^5 = \{b_0 = \text{None} = N, b_1 = \text{Low} = L, b_2 = \text{Medium} = M, b_3 = \text{High} = H, b_4 = \text{Maximum} = MA\}$
- $S^9 = \{c_0 = \text{None} = N, c_1 = \text{Very_Low} = VL, c_2 = \text{Low} = L, c_3 = \text{Slightly_Low} = SL, c_4 = \text{Medium} = M, c_5 = \text{Slightly_High} = SH, c_6 = \text{High} = H, c_7 = \text{Very_High} = VH, c_8 = \text{Maximum} = MA\}$

In [13] authors remark that the family of transformation functions between labels from different levels is bijective, guarantying that the transformations between levels are produced without loss of information in a linguistic hierarchy.

3 Proposal Description

In this section we present an automatic academic profiles builder for users in order to deal with the cold start problem using multi-granular fuzzy linguistic modeling and bibliometrics. We work over REFORE, a recommender system for researchers introduced in our previous work [24]. First, we will see the architecture and approach followed. Then, we will go through the representation of the information as well as the resources. We will conclude with profile formation of the researchers.

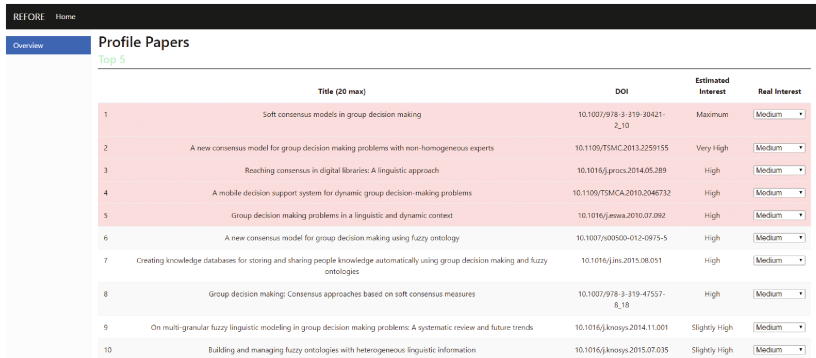
3.1 System Concepts

The approach works based on the following concepts:

- Researchers does not want to spend time building a profile. However there is a better predisposition to adjust one done beforehand.
- *Cold Star* problem is solved adding extra information from users from additional data sources. The system extracts all the necessary information from Scopus based on the name of the researcher.
- Due to the nature of research itself, authors needs of information or investigation interests are very specific.
- The use of bibliometric quality measures worked in REFORE [24] as part of the recommendation approach. Authors tend to public their best works in the best possible journals.

- Novelty. To follow the hype. Authors will be more interested in topics related with their last works.
- Frequency. The more recurrent topics from authors will mark an important lines
- Authorship. Authors tend to be first authors in those works they lead.

REFORE profiles key needs for its correct functioning are two: keywords for the main search and papers for the bibliometrics filters applied. The process followed by the profile builder after a researcher name is introduced is shown on Fig. 2 and consists on:



The screenshot shows the REFORE application interface. At the top, there is a navigation bar with 'REFORE' and 'Home'. Below it, a sidebar contains 'Character' and 'Profile Papers'. The 'Profile Papers' section is active, displaying 'Top 5' papers. The table below lists these papers with their titles, DOIs, estimated interest, and real interest.

	Title (20 max)	DOI	Estimated Interest	Real Interest
1	Soft consensus models in group decision making	10.1007/978-3-319-30421-2_10	Maximum	Medium
2	A new consensus model for group decision making problems with non-homogeneous experts	10.1030/TSMC.2013.2259155	Very High	Medium
3	Reaching consensus in digital libraries: A linguistic approach	10.1016/j.procs.2014.05.269	High	Medium
4	A mobile decision support system for dynamic group decision-making problems	10.1109/TSMA.2010.2046732	High	Medium
5	Group decision making problems in a linguistic and dynamic context	10.1016/j.asmea.2010.07.082	High	Medium
6	A new consensus model for group decision making using fuzzy ontology	10.1007/s00500-012-0975-5	High	Medium
7	Creating knowledge databases for storing and sharing people knowledge automatically using group decision making and fuzzy ontologies	10.1016/j.jas.2015.08.051	High	Medium
8	Group decision making: Consensus approaches based on soft consensus measures	10.1007/978-3-319-47557-8_19	High	Medium
9	On multi-granular fuzzy linguistic modeling in group decision making problems: A systematic review and future trends	10.1016/j.knsys.2014.11.001	Slightly High	Medium
10	Building and managing fuzzy ontologies with heterogeneous linguistic information	10.1016/j.knsys.2015.07.035	Slightly High	Medium

Fig. 2. Operating scheme

1. Query the Scopus API ¹ for author information being papers the most important.
2. Estimate the importance degree of the paper on the author profile. Details explained below.
3. Extracting keywords from all the user papers weighted with the importance of each.

3.2 Information Representation

In order to represent the information we stayed with the same linguistic hierarchy defined in REFORE [24] but using only two levels. The concepts to asses in this work are the following:

- The *Importance degree* of keywords for the users, which is assessed in S^5 .
- The *Relevance degree* of a paper for a user, which is assessed in S^9 .

We propose to use a linguistic hierarchy which linguistics term sets are:

¹ <https://dev.elsevier.com/scopus.html>.

- $S^5 = \{b_0 = \text{None} = N, b_1 = \text{Low} = L, b_2 = \text{Medium} = M, b_3 = \text{High} = H, b_4 = \text{Maximum} = MA\}$
- $S^9 = \{c_0 = \text{None} = N, c_1 = \text{VeryLow} = VL, c_2 = \text{Low} = L, c_3 = \text{SlightlyLow} = SL, c_4 = \text{Medium} = M, c_5 = \text{SlightlyHigh} = SH, c_6 = \text{High} = H, c_7 = \text{VeryHigh} = VH, c_8 = \text{Maximum} = MA\}$

Level 1 is used to represent the importance degree of keywords ($S_1 = S^9$) and for the predicted relevance degrees we use the level 2 ($S_2 = S^5$).

3.3 Profiles Construction

Author names are required in the format used for publishing. Afterwards, Scopus API is used to look for authors profiles and retrieve them together with their research history. The system performs a quick analysis on the papers splitting user profiles in two: *Papers* and *Keywords*.

In REFORE, *papers* were split in two groups, *selected* and *non-selected*. This classification was used for the filtering process. In order to provide the same classification, the following characteristics have been considered to estimate the individual importance degree of each one.

Given a paper P , P_{iu} is the paper i from user u estimated aggregating the following paper characteristics related to the user:

- Quality, given by the *Impact Factor*:

$$IF_{iu} \begin{cases} IF_{i,J} & \text{if } J = \text{ranked journal} \\ 0.4 & \text{Otherwise} \end{cases} \quad (1)$$

where J is the source where the paper i from user u was published. We consider authors tend to publish their better works in the best journals, giving slightly less importance to conferences.

- *Novelty*, given by the publication date:

$$No_{iu} \begin{cases} \text{Current_Year} - \text{Year}(P_{iu}) * 0.5 & \text{if } \text{Current_Year} - Y > 5 \\ 0.4 & \text{Otherwise} \end{cases} \quad (2)$$

where $\text{Year}(P_{iu})$ is the year when the paper P_{iu} was published. We considered that authors tend to be more interested on the research lines they are working in the present.

- *Authorship*, given by the occupied position on the authors line in the paper:

$$Au_{iu} \begin{cases} 2 & \text{if } \text{Current_Year} - Y > 5 \\ 0.4 & \text{Otherwise} \end{cases} \quad (3)$$

where being the first author in the list is considered as being the lead carrier of work.

Before aggregating, all results are normalized within the interval $[0, 1]$. Different weights distribution are applied: 20% for the IF , 50% for the No and 30% the Au .

Given the necessity of operating in the following steps with the linguistic set of keywords which are expressed with labels from level 2 of our linguistic hierarchy, that is $S_1 = S^9$, a linguistic transformation is needed. We followed the representation model described in [12] to transform the aggregated values to their linguistic labels belonging to the level 3 of our linguistic hierarchy, i.e.: $S_1 = S^9$. Level 3 was chosen to allow users a bigger margin when reviewing the papers profile.

On the other hand, *keywords* are used for the similarity estimation in REFORE. In order to obtain the importance degree on each of them we applied the linguistic weighted average (see Definition 1). The importance degree results as the average of each keyword appearing on user papers, each keyword inherits the importance degree of the paper it appears.

Definition 1. Linguistic Weighted Average Operator [24]. Let $x = \{(r_1, \alpha_1), \dots, (r_n, \alpha_n)\}$ be a set of linguistic 2-tuples and $W = \{(w_1, \alpha_1^w), \dots, (w_n, \alpha_n^w)\}$ be their linguistic 2-tuple associated weights. The 2-tuple linguistic weighted average $\bar{x}_l^w$ is:

$$\bar{x}_l^w [((r_1, \alpha_1), (w_1, \alpha_1^w)) \dots ((r_n, \alpha_n), (w_n, \alpha_n^w))] = \Delta \left(\frac{\sum_{i=1}^n \beta_i \cdot \beta_{W_i}}{\sum_{i=1}^n \beta_{W_i}} \right), \quad (4)$$

with $\beta_i = \Delta^{-1}(r_i, \alpha_i)$ and $\beta_{W_i} = \Delta^{-1}(w_i, \alpha_i^w)$.

As above mentioned keywords are expressed withing the level 2 of the linguistic hierarchy used. So in order to be used by REFORE a transformation between level is performed.

4 Experiments and Approach Evaluation

In this section we present the evaluation of the proposed approach. Due to the nature of REFORE, the system objective of this approach proposition is to alleviate the cold start problem. No comparison with other approaches is possible since no standard data set is used. Thus, in this study we only perform online experiments, i.e., practical studies where a group of researchers indicate their optimum profile. Users input is compared to the estimated one by our approach.

In order to test the effectiveness of the approach followed and after adjusting the different parameters to the optimum weight, i.e.: Quality, novelty and authorship, the experiment is performed over a set of users from REFORE. Other approaches test the validity of their solutions for cold start problem through evaluating recommendations, in this work we propose a direct evaluation of the profiles by the own users. So our experiment consists on showing to the users the profiles built for them together with the estimated importance values for each paper and keyword. Afterwards, the user is inquired for the real ones.

For that purpose a section in REFORE was created for the test group (see Fig. 3).

The test group consisted on 20 researchers with different profiles, going from a more junior research profiles to senior ones. After loading their profiles from

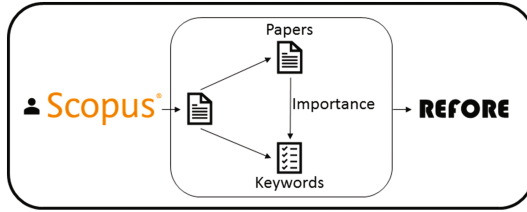


Fig. 3. Evaluation page

Scopus extracting papers and keywords, the system estimates their importance degree following the steps indicated in Sect. 3 and shows them to the user.

When it comes to recommender systems, the most common measures for accuracy are precision, recall or F1 [20]. However, in this work we left to users the evaluation of their own estimated profile, so in order to measure how accurate the system is mirroring user interests we use the **Mean Absolute Error (MAE)** [10]. In particular, we defined it in a linguistic framework:

$$MAE = \Delta(g \times \frac{\sum_{i=1}^n abs(\Delta^{-1}(p_i, \alpha_{pi}) - \Delta^{-1}(r_i, \alpha_{ri}))}{n}) \quad (5)$$

with $MAE \in S_1 \times [-0.5, 0.5]$, and where g is the granularity of the term set used to express the relevance degree, i.e. S^g , n is the the number of cases in the test set, (p_i, α_{pi}) is the predicted 2-tuple linguistic value for paper or keyword i and (r_i, α_{ri}) the real one.

We evaluated both, papers and keywords separately, importance degree for keywords and relevance degree for papers, obtaining the following MAE results:

- **Keywords:** 0.127
- **Papers:** 0.104

We observe that the profiles estimations generated with the proposed approach are in line with the real users preferences, softening or removing the necessity of their interaction with the system to establish a previous user profile.

5 Concluding Remarks

In this paper we propose a fuzzy linguistic approach based on bibliometrics to deal with the cold start problem for researchers present in the REFORE system.

As we experienced in REFORE, the first barrier users find to use a recommender system is the building of their own profiles. The idea of automatize the profile construction will support the system as well as will enable the creation of different system more oriented to real time interactions.

A user historic research record is retrieved from Scopus. The system split keywords and the rest of meta-information from the papers from authors. Afterwards, the importance degree of each paper and keyword is obtained based on

the idea of novelty, frequency and quality. We have applied the approach over our previous recommender system REFORE in a real environment with satisfactory results. Those results, showed that the approach performance was better within the top 5 element of the each list, keywords and papers.

As future work, we consider to study the inclusion of the automatic profiling on a real time recommender system for researcher with no need of previous registration.

Acknowledgments. This paper has been developed with the FEDER financing under Projects TIN2013-40658-P and TIN2016-75850-R

References

1. Bordogna, G., Pasi, G.: A fuzzy linguistic approach generalizing boolean information retrieval: A model and its evaluation. *J. Am. Soc. Inf. Sci.* **44**(2), 70 (1993)
2. Burke, R.: Hybrid web recommender systems. In: Brusilovsky, P., Kobsa, A., Nejdl, W. (eds.) *The Adaptive Web*. LNCS, vol. 4321, pp. 377–408 (2007)
3. Chen, P.Y., Wu, S.Y.: Does collaborative filtering technology impact sales? empirical evidence from amazon.com (2007). <https://ssrn.com/abstract=1002698>
4. Claypool, M., Gokhale, A., Miranda, T.: Combining content-based and collaborative filters in an online newspaper. In: *Proceedings of the ACM SIGIR 1999 Workshop on Recommender Systems-Implementation and Evaluation*, pp. 714–720 (1999)
5. Delgado, M., Herrera, F., Herrera-Viedma, E., Martínez, L.: Combining numerical and linguistic information in group decision making. *Inf. Sci.* **107**(1–4), 177–194 (1998)
6. Edmunds, A., Morris, A.: The problem of information overload in business organizations: a review of the literature. *Int. J. Inf. Manage.* **20**, 17–28 (2000)
7. Esteban, B., Tejeda-Lorente, A., Porcel, C., Arroyo, M., Herrera-Viedma, E.: TPLUFIB-WEB: A fuzzy linguistic web system to help in the treatment of low back pain problems. *Knowl. Based Syst.* **67**, 429–438 (2014)
8. Good, N., Schafer, J., Konstan, J., Borchers, A., Sarwar, B., Herlocker, J., Riedl, J.: Combining collaborative filtering with personal agents for better recommendations. In: *Proceedings of the Sixteenth National Conference on Artificial Intelligence, AAAI/IAAI*, pp. 439–446 (1999)
9. Hanani, U., Shapira, B., Shoval, P.: Information filtering: Overview of issues, research and systems. *User Model. User-Adap. Inter.* **11**, 203–259 (2001)
10. Herlocker, J., Konstan, J., Terveen, L., Riedl, J.: Evaluating collaborative filtering recommender systems. *ACM Trans. Inf. Syst.* **22**(1), 5–53 (2004)
11. Herrera, F., Herrera-Viedma, E.: Choice functions and mechanisms for linguistic preference relations. *Eur. J. Oper. Res.* **120**, 144–161 (2000)
12. Herrera, F., Martínez, L.: A 2-tuple fuzzy linguistic representation model for computing with words. *IEEE Trans. Fuzzy Syst.* **8**(6), 746–752 (2000)
13. Herrera, F., Martínez, L.: A model based on linguistic 2-tuples for dealing with multigranularity hierarchical linguistic contexts in multiexpert decision-making. *IEEE Trans. Syst. Man Cybern. Part B Cybern.* **31**(2), 227–234 (2001)
14. Herrera-Viedma, E., López-Herrera, A., Luque, M., Porcel, C.: A fuzzy linguistic model based on a 2-tuple fuzzy linguistic approach. *Int. J. Uncertainty Fuzziness Knowl. Based Syst.* **15**(2), 225–250 (2007)

15. Hsu, M.: A personalized english learning recommender system for ESL students. *Expert Syst. Appl.* **34**, 377–408 (2008)
16. Kotsiantis, S.B.: Supervised machine learning: A review of classification techniques. In: *Proceedings of the 2007 Conference on Emerging Artificial Intelligence Applications in Computer Engineering: Real Word AI Systems with Applications in eHealth, HCI, Information Retrieval and Pervasive Technologies*, pp. 3–24. IOS Press, Amsterdam, The Netherlands (2007)
17. Lika, B., Kolomvatsos, K., Hadjiefthymiades, S.: Facing the cold start problem in recommender systems. *Expert Syst. Appl.* **41**(4, Part 2), 2065–2073 (2014)
18. Liu, H., Hu, Z., Mian, A., Tian, H., Zhu, X.: A new user similarity model to improve the accuracy of collaborative filtering. *Knowl. Based Syst.* **56**, 156–166 (2014)
19. Long-Sheng, C., Fei-Hao, H., Mu-Chen, C., Yuan-Chia, H.: Developing recommender systems with the consideration of product profitability for sellers. *Inf. Sci.* **178**(4), 1032–1048 (2008)
20. Lops, P., Gemmis, M., Semeraro, G.: Content-based recommender systems: State of the art and trends. In: Ricci, F., Rokach, L., Shapira, B., Kantor, P.B. (eds.) *Recommender Systems Handbook*, pp. 73–106. Springer (2011)
21. Rish, I.: An empirical study of the naive bayes classifier. In: *IJCAI 2001 Workshop on Empirical Methods in Artificial Intelligence*, Vol. 3, pp. 41–46. IBM New York (2001)
22. Schein, A.I., Popescul, A., Ungar, L.H., Pennock, D.M.: Methods and metrics for cold-start recommendations. In: *Proceedings of the 25th annual international ACM SIGIR conference on Research and development in information retrieval*, pp. 253–260. ACM (2002)
23. Son, L.H.: Dealing with the new user cold-start problem in recommender systems: A comparative review. *Inf. Syst.* **58**, 87–104 (2016)
24. Tejeda-Lorente, A., Porcel, C., Bernabé-Moreno, J., Herrera-Viedma, E.: Refore: A recommender system for researchers based on bibliometrics. *Appl. Soft Comput.* **30**, 778–791 (2015)
25. Tejeda-Lorente, A., Porcel, C., Peis, E., Sanz, R., Herrera-Viedma, E.: A quality based recommender system to disseminate information in a university digital library. *Inf. Sci.* **261**, 52–69 (2014)
26. Zadeh, L.: The concept of a linguistic variable and its applications to approximate reasoning. Part I, *Inf. Sci.* 8, 199–249 (1975), Part II, *Inf. Sci.* 8, 301–357 (1975), Part III. *Inf. Sci.* 9(1975), 43–80 (1975)