



A hybrid recommender system for the selective dissemination of research resources in a Technology Transfer Office

C. Porcel^{a,*}, A. Tejeda-Lorente^b, M.A. Martínez^b, E. Herrera-Viedma^{b,*}

^a University of Jaén, Department of Computer Science, Jaén, Spain

^b University of Granada, Department of Computer Science and Artificial Intelligence, Granada, Spain

ARTICLE INFO

Article history:

Received 14 December 2009

Received in revised form 6 June 2011

Accepted 22 August 2011

Available online 6 September 2011

Keywords:

Recommender systems

Fuzzy linguistic modeling

Technology Transfer Office

ABSTRACT

Recommender systems could be used to help users in their access processes to relevant information. Hybrid recommender systems represent a promising solution for multiple applications. In this paper we propose a hybrid fuzzy linguistic recommender system to help the Technology Transfer Office staff in the dissemination of research resources interesting for the users. The system recommends users both specialized and complementary research resources and additionally, it discovers potential collaboration possibilities in order to form multidisciplinary working groups. Thus, this system becomes an application that can be used to help the Technology Transfer Office staff to selectively disseminate the research knowledge and to increase its information discovering properties and personalization capacities in an academic environment.

© 2011 Elsevier Inc. All rights reserved.

1. Introduction

Theoretical and empirical works in innovation economics suggest that the use of scientific knowledge by setting up and maintaining good industry/science relations positively affects innovation performance [48]. In terms of organizational structure, creating a specialized Technology Transfer Office (TTO) within a university can be instrumental in developing relations with the industry. To arrange a dedicated transfer unit, that acts as “technological intermediaries”, allows for specialization in support services, most notably, partner search, management of intellectual property, and business development [48,50]. The Technology Transfer Offices were established to facilitate commercial knowledge transfers from universities to practitioners or university/industry technology transfer [69]. They are responsible for managing and putting into action the activities which generate knowledge and technical and scientific collaboration, thus enhancing the interrelation between researchers at the university and the entrepreneurial world and their participation in various support programmes designed to carry out research, development and innovation activities.

A TTO develops a range of services to carry out its objectives:

- Guidance for Research and Development (R&D) and Technology Transfer funding.
- Disseminate information (R&D bulletins, R&D&I, calls, notices, projects and so on).
- Advice in the preparation of offers (management, spread and exploitation).
- Support in the elaboration and negotiation of contracts with companies.
- Management of contacts.

* Corresponding authors.

E-mail addresses: cporcel@ujaen.es (C. Porcel), atejeda@decsai.ugr.es (A. Tejeda-Lorente), mnglsmartinez@gmail.com (M.A. Martínez), viedma@decsai.ugr.es (E. Herrera-Viedma).

- Technological offers (the elaboration of the offer, spread and promotion).
- The advice in the creation of new businesses.
- Evaluation, protection and transfer of both intellectual and industrial ownership rights.

However, a TTO is a separate unit, and therefore it needs to maintain close relationships with researchers in the different departments and have the proper incentive mechanisms in place to ensure researchers to generate inventions and disclose them to the TTO [48]. In this sense, a service that is particularly important to fulfill this objective is the selective dissemination of information about research resources. It would allow to increase the visibility of the academic departments and research groups to the society and industry. Empirical evidence shows that the interactions between universities and industry have intensified in recent years and therefore the number of available resources is growing quickly [48]. So the TTO staff finds difficulties in achieving an effective selective dissemination of information. To solve this problem, automatic techniques are needed in the TTO to facilitate users to selectively access to research resources. Mainly, there exist two different tools to facilitate the access to the information: Information Retrieval Systems [34,44,47] and Recommender Systems [4,20,45,53,58,60,63,70]. The former are focused on information search in a known content repository while the later are focused on information discovery in partially known frameworks. A recommender system attempts to discover information items (movies, music, books, news, images, web pages, papers and so on) that are likely of interest to a user. Recommender systems are especially useful when they identify information that a person was previously unaware of. Furthermore, recommender systems are personalized services because they may treat each user in a different way. From a theoretical point of view, recommender systems have fallen into two main categories [19,20,22,54,57,63,67,70]:

1. *Content-based recommender systems* recommend information items to a user by means of a process based on the content of the information item and the user's past experience dealing with similar items, and therefore, ignoring data from other users.
2. *Collaborative recommender systems* recommend information items to a user by means of a process based on the user's social environment and ignoring the contents of the items, that is, the recommendations to a user are based on other user recommendations with similar user profiles.

On the other hand, from a practical point of view, four different types of recommendations are identified [40]:

1. *Personalized recommendations* recommend items based on the individual's past behavior, as in the content-based filtering.
2. *Social recommendations* recommend items based on the past behavior of similar users, as in the collaborative filtering.
3. *Item recommendations* recommend items based on the item itself, as it happens in information retrieval systems [34,44,47] but assuming long time queries.
4. A combination of the three approaches described above.

In [59] we presented SIRE2IN, a fuzzy recommender system to help the TTO staff of the University of Granada in the management of research resources, as calls for research projects and so on. This system uses a fuzzy linguistic modeling to represent the qualitative information presented in the system communication processes [6,12,13,23,26,27,42,72–74]. Particularly, we use a multi-granular fuzzy linguistic modeling [28,37,31] that provides greater flexibility in the user-system interaction, which turns to be an interesting and useful characteristic. However, we have found different aspects that may limit its performance:

1. It acts as an information retrieval system based on matching functions which acts among the resources representation and user profiles and therefore, it does not implement a recommender system core to discover new information to the users. Furthermore, it only shares user experiences or social wisdom [66].
2. The system looks for community members to collaborate in a limited way because it recommend researchers who present similar profiles to the user. Then, the system obtains accurate but unhelpful recommendations.
3. TTO staff have found serious troubles to achieve an effective customization in the information dissemination processes with SIRE2IN.
4. Nowadays, as it happens in the Web, the TTO suffers the information overload problem. The number of electronic resources daily generated grows and we have found that the SIRE2IN performance has decreased.

If we analyze the TTO scope, we find that the collaborative filtering approach is very useful because it allows users to share their experiences, that is, users can rate or add value to research resources and these ratings can be shared with the community, so that popular resources can be easily located or people can receive information items found useful by others with similar profiles. Consequently and, taking into account the difficulties found in the previous system, in this paper we replace the recommendation scheme by a hybrid approach.

The aim of this paper is to present a new fuzzy linguistic recommender system which is applied in the TTO in the University of Granada. This new system implements a hybrid recommendation approach which improves the SIRE2IN performance with respect to the discovering and personalization capacities. In such a way, it allows to help the TTO staff to

selectively disseminate research knowledge and the researchers to discover information. The most important novelties of this new fuzzy linguistic recommender system are:

- The system implements a hybrid recommendation strategy based in a *switching hybrid approach* [7], which switches between a content-based recommendation approach and a collaborative one to share the user individual experience and social wisdom.
- The system implements a personalization tool that allows to recommend users three types of items:
 1. Specialized resources of the own user research area to contribute to his/her specialization.
 2. Other resources as complementary formation.
 3. Research collaborators. In this case, it allows researchers to discover new members with complementary profiles which could provide them real collaboration possibilities to form multidisciplinary working groups and develop common projects.
- The system implements a richer feedback process. In [59] the user participation in the recommendation process is small because the user feedback consists of adding or eliminating topics in the user profile, but users could not provide satisfaction degrees. However, to improve the recommendations in this new system, when researchers analyze a recommended resource they provide a satisfaction degree. In such a way, we guarantee that user experiences are taken into account to generate the recommendations done by the system.

The system has been developed in the University of Granada and the experimental results show us that it is useful and effective for the users. Besides, in order to compare the results of our system with other, we have implemented several content-based and collaborative models. The results reflect an improvement in the system performance when we use the new recommendation kernel.

The paper is structured as follows. Section 2 presents the basic concepts and aspects about the recommender systems. Section 3 revises the multi-granular fuzzy linguistic modeling. In Section 4 we present the new recommender system to selectively advice research resources in a TTO. Section 5 reports the system evaluation and the experimental results. Finally, our concluding remarks are pointed out in Section 6.

2. Basis of recommender systems

Recommender systems help online users in the effective identification of items suiting their wishes, needs or preferences. They have the effect of guiding the users in a personalized way to relevant or useful objects in a large space of possible options [7]. These applications improve the information access processes for users not having a detailed product domain knowledge. They are becoming popular tools for reducing information overload and to improve the sales in e-commerce web sites [8,11,17,18,41,46,63].

Automatic filtering services differ from retrieval services. In filtering the corpus changes continuously, the users have long time information needs (described by mean of user profiles), and the objective is to remove irrelevant data from incoming streams of data items [17,20,49,63]. On the contrary, retrieval services use queries which are introduced by the users into the system to retrieve relevant items. Thus, a result from a recommender system is understood as a recommendation, an option worthy or consideration, while a result from an information retrieval system is interpreted as a match to the user's query [8].

In a recommender system, the users' preferences about research resources can be used to define user profiles that are applied as filters to streams of documents. The construction of accurate profiles is a key task and the system's success will depend on a large extent on the ability of the learned profiles to represent the user's preferences [61]. Then, in order to generate personalized recommendations that are tailored to the user's preferences or needs, recommender systems must collect personal preference information, such as user's history of purchase, items which were previously interesting for the user, click-stream data, demographic information, and so on.

Two different ways to obtain information about user preferences are distinguished [20], although many systems adopt a hybrid approach:

- The *implicit approach* is implemented by inference from some kind of observation. The observation is applied to user behavior or to detect a user's environment (such as bookmarks or visited URL). The user preferences are updated by detecting changes while observing the user.
- The *explicit approach* interacts with the users by acquiring feedback on information that is filtered, that is, the users express some specifications of what they desire. This approach is currently the most common one.

In addition, there are mainly two approaches that have been proposed to implement recommender applications [19,20,54,57,63,67,70]:

- *Content-based systems*: They generate the recommendations taking into account the characteristics used to represent the items and the ratings that a user has given to them [5,14]. These recommender systems tend to fail when little is known about the user information needs. This is called the new user cold-starting problem [43].

- *Collaborative systems*: The system generates recommendations using explicit or implicit preferences from many users, ignoring the items representation. Collaborative systems locate peer users with a rating history similar to the current user and they generate recommendations using this neighborhood. These recommender systems tend to fail when little is known about items, i.e., when new items appear. This is called the new item cold-starting problem [8].

In this paper we propose the use of a hybrid approach to reduce the disadvantages of each one of them and to exploit their benefits. Using a hybrid strategy users are provided with more accurate recommendations than those offered by each strategy individually [5,14,19].

Usually, the recommendation activity is followed by a relevance feedback phase. Relevance feedback is a cyclic process whereby the users provide the system with their satisfaction evaluations about the recommended items and the system uses these evaluations to automatically update user profiles in order to generate new recommendations [20,63].

3. Multi-granular fuzzy linguistic modeling

In this section we present the multi-granular fuzzy linguistic approach used in our recommender system.

The use of Fuzzy Sets Theory has given very good results to model qualitative information [73] and it has been proven to be useful in many problems, e.g., decision making [2,10,24,26,42,75], quality evaluation [9,38,39,55], information retrieval [29,30,32,33,35,36], political analysis [3], estimation of student performances [56], etc. It is a tool based on the concept of *linguistic variable* proposed by Zadeh [73].

In any fuzzy linguistic approach, an important parameter to determine is the granularity of uncertainty, i.e., the cardinality of the linguistic term set S . According to the uncertainty degree that an expert qualifying a phenomenon has on it, the linguistic term set chosen to provide his knowledge will have more or less terms. When different experts have different uncertainty degrees on the phenomenon, then several linguistic term sets with a different granularity of uncertainty are necessary [25]. The use of different label sets to assess information is also necessary when an expert has to evaluate different concepts, as it happens in information retrieval problems when users have to evaluate the importance of the query terms and the relevance of the retrieved documents [31]. In such situations, we need tools to manage multi-granular linguistic information [28,37,51].

In [28] a multi-granular fuzzy linguistic modeling based on a 2-tuple fuzzy linguistic approach and the concept of linguistic hierarchy was proposed.

3.1. The 2-tuple fuzzy linguistic approach

The 2-tuple fuzzy linguistic modeling [27] is a continuous model of information representation that allows to reduce the loss of information that typically arise when using other fuzzy linguistic approaches (classical and ordinal [23,73]). To define it both the 2-tuple representation model and the 2-tuple computational model to represent and aggregate the linguistic information have to be established.

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 symmetric related to it. We assume that the semantics of labels is given by means of fuzzy subsets defined in the $[0, 1]$ interval, which are described by their membership functions $\mu_{s_i} : [0, 1] \rightarrow [0, 1]$, and we consider all terms distributed on a scale on which a total order is defined, that is, $s_i \leq s_j \rightarrow i \leq j$. We consider that linear triangular membership functions are good enough to capture the vagueness of those linguistic assessments, since it may be impossible or unnecessary to obtain more accurate values. This representation is achieved by the 3-tuple (a, b, c) , where a is the point where the membership is 1 and b and c are the left and right limits of the definition domain of the triangular membership function. For example, the following semantics, represented in Fig. 1, can be assigned to a set of seven terms via triangular membership functions:

$$\begin{aligned} P = \text{Perfect} &= (0.83, 1, 1) \quad VH = \text{Very_High} = (0.67, 0.83, 1), \\ H = \text{High} &= (0.5, 0.67, 0.83) \quad M = \text{Medium} = (0.33, 0.5, 0.67), \\ L = \text{Low} &= (0.17, 0.33, 0.5) \quad VL = \text{Very_Low} = (0, 0.17, 0.33), \\ N = \text{None} &= (0, 0, 0.17). \end{aligned}$$

In this fuzzy linguistic context, if a symbolic method [23,26] 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 .

Definition 1 [27]. Let β be the result of an aggregation of the indexes of a set of labels assessed in a linguistic term set S , i.e., the result of a symbolic aggregation operation, $\beta \in [0, g]$. Let $i = \text{round}(\beta)$ and $\alpha = \beta - i$ be two values, such that, $i \in [0, g]$ and $\alpha \in [-.5, .5]$ then:

- s_i represents the linguistic label of the information, and
- α_i is a numerical value expressing the value of the symbolic translation from the original result β to the closest index label, i , in the linguistic term set ($s_i \in S$).

This model defines a set of transformation functions between numeric values and 2-tuples.

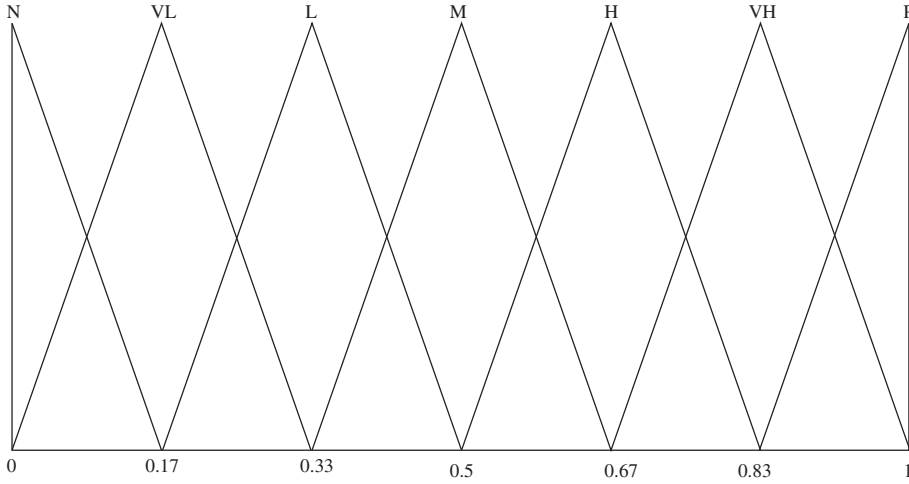


Fig. 1. A set of seven linguistic terms with its semantics.

Definition 2 [27]. Let $S = \{s_0, \dots, s_g\}$ be a linguistic term set and $\beta \in [0, g]$ a value representing the result of a symbolic aggregation operation, then the 2-tuple that expresses the equivalent information to β is obtained with the following function:

$$\Delta : [0, g] \rightarrow S \times [-0.5, 0.5),$$

$$\Delta(\beta) = (s_i, \alpha), \text{ with } \begin{cases} s_i & i = \text{round}(\beta), \\ \alpha = \beta - i & \alpha \in [-.5, .5), \end{cases}$$

where $\text{round}(\cdot)$ is the usual *round* operation, s_i has the closest index label to “ β ” and “ α ” is the value of the symbolic translation.

For all Δ there exists Δ^{-1} , defined as $\Delta^{-1}(s_i, \alpha) = i + \alpha$. On the other hand, it is obvious that the conversion of a linguistic term into a linguistic 2-tuple consists of adding a symbolic translation value of 0: $s_i \in S \Rightarrow (s_i, 0)$.

The computational model is defined by presenting the following operators:

1. Negation operator: $\text{Neg}((s_i, \alpha)) = \Delta(g - (\Delta^{-1}(s_i, \alpha)))$.
2. Comparison of 2-tuples (s_k, α_1) and (s_l, α_2) :
 - If $k < l$ then (s_k, α_1) is smaller than (s_l, α_2) .
 - If $k = l$ then
 - (a) if $\alpha_1 = \alpha_2$ then (s_k, α_1) and (s_l, α_2) represent the same information,
 - (b) if $\alpha_1 < \alpha_2$ then (s_k, α_1) is smaller than (s_l, α_2) ,
 - (c) if $\alpha_1 > \alpha_2$ then (s_k, α_1) is bigger than (s_l, α_2) .
3. Aggregation operators [71,74]. The aggregation of information consists of obtaining a value that summarizes a set of values, therefore, the result of the aggregation of a set of 2-tuples must be a 2-tuple. In the literature we can find many aggregation operators which allow us to combine the information according to different criteria. Using functions Δ and Δ^{-1} that transform without loss of information numerical values into linguistic 2-tuples and viceversa, any of the existing aggregation operators can be easily extended to deal with linguistic 2-tuples. Some examples are:

Definition 3 (Arithmetic Mean). Let $x = \{(r_1, \alpha_1), \dots, (r_n, \alpha_n)\}$ be a set of linguistic 2-tuples, the 2-tuple arithmetic mean $\bar{x}^e$ is computed as,

$$\bar{x}^e[(r_1, \alpha_1), \dots, (r_n, \alpha_n)] = \Delta\left(\sum_{i=1}^n \frac{1}{n} \Delta^{-1}(r_i, \alpha_i)\right) = \Delta\left(\frac{1}{n} \sum_{i=1}^n \beta_i\right).$$

Definition 4 (Weighted Average Operator). Let $x = \{(r_1, \alpha_1), \dots, (r_n, \alpha_n)\}$ be a set of linguistic 2-tuples and $W = \{w_1, \dots, w_n\}$ be their associated weights. The 2-tuple weighted average $\bar{x}^w$ is:

$$\bar{x}^w[(r_1, \alpha_1), \dots, (r_n, \alpha_n)] = \Delta\left(\frac{\sum_{i=1}^n \Delta^{-1}(r_i, \alpha_i) \cdot w_i}{\sum_{i=1}^n w_i}\right) = \Delta\left(\frac{\sum_{i=1}^n \beta_i \cdot w_i}{\sum_{i=1}^n w_i}\right).$$

Definition 5 (*Linguistic Weighted Average Operator*). 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:

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

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

3.2. Linguistic hierarchy to model multi-granular linguistic information

A *Linguistic Hierarchy, LH*, is a set of levels $l(t, n(t))$, i.e., $LH = \bigcup l(t, n(t))$, where each level t is a linguistic term set with different granularity $n(t)$ from the remaining of levels of the hierarchy. 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: $l(t, n(t)) \rightarrow l(t + 1, 2 \cdot n(t) - 1)$. Table 1 shows the granularity needed in each linguistic term set of the level t depending on the value $n(t)$ defined in the first level (3 and 7 respectively).

A graphical example of a linguistic hierarchy is shown in Fig. 2.

In [28] was demonstrated that linguistic hierarchies are useful to represent multi-granular linguistic information and that they allow to combine multi-granular linguistic information without loss of information. To do this, a family of transformation functions between labels from different levels was defined:

Definition 6. Let $LH = \bigcup l(t, n(t))$ be a linguistic hierarchy whose linguistic term sets are denoted as $S^{n(t)} = \{s_0^{n(t)}, \dots, s_{n(t)-1}^{n(t)}\}$. The transformation function between a 2-tuple that belongs to level t and another 2-tuple in level $t' \neq t$ is defined as:

$$TF_{t'}^t : l(t, n(t)) \rightarrow l(t', n(t')),$$

$$TF_{t'}^t \left(s_i^{n(t)}, \alpha^{n(t)} \right) = \Delta \left(\frac{\Delta^{-1} \left(s_i^{n(t)}, \alpha^{n(t)} \right) \cdot (n(t') - 1)}{n(t) - 1} \right).$$

As it was pointed out in [28] this family of transformation functions is bijective. This result guarantees the transformations between levels of a linguistic hierarchy are carried out without loss of information. To define the computational model, we select a level to make the information uniform (for instance, the highest granularity level) and then we can use the operators defined in the 2-tuple fuzzy linguistic approach.

Table 1
Linguistic hierarchies.

	Level 1	Level 2	Level 3
$l(t, n(t))$	$l(1, 3)$	$l(2, 5)$	$l(3, 9)$
$l(t, n(t))$	$l(1, 7)$	$l(2, 13)$	

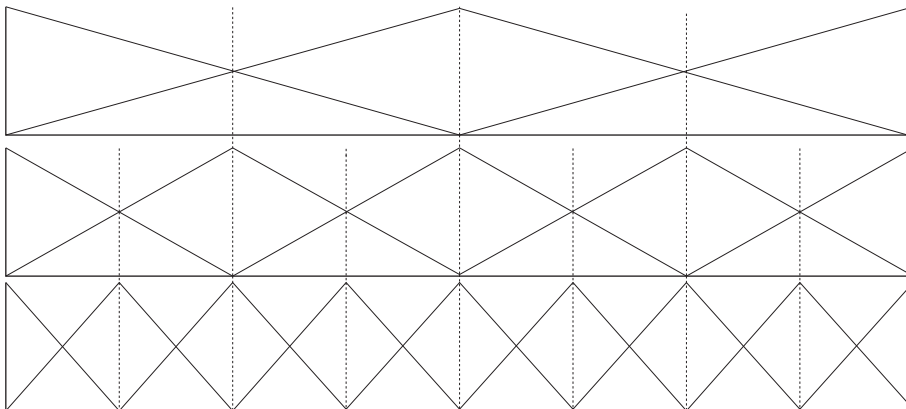


Fig. 2. Linguistic Hierarchy of 3, 5 and 9 labels.

4. A recommender system for the selective dissemination of research resources in a TTO

In this section, we present a new fuzzy linguistic hybrid recommender system to discover both researchers information about research resources and collaboration possibilities. The TTO staff manages and spreads knowledge about research resources such as R&D bulletins, R&D&I, calls, notices, research projects and so on [48,50]. Nowadays, this amount of information grows continuously and the TTO staff needs automated tools to filter and spread that information to the researchers in a simple and timely manner.

As aforementioned, the aim of this new system is to overcome some of the problems detected in SIRE2IN, and so, we present an improved system which can be used in a real TTO environment to achieve an effective selective dissemination of research resources. This system works according to a hybrid recommendation strategy based in a switching hybrid approach [7], which switches between a content-based recommendation approach and a collaborative one to share user experiences by generating social recommendations. Basically, the former is applied when a new item is inserted and the latter is applied when a new researcher is registered. Furthermore, we include another novelty which suggests resources to researchers, recommending both specialized and complementary research resources. It also improves the recommendation process, allowing researchers to discover real collaboration possibilities in order to form multidisciplinary working groups. In such a way, the new system improves the services of a TTO, selectively disseminating research resources, and allowing to share knowledge in an academic context.

We present a multi-granular fuzzy linguistic recommender system that provides high flexibility in the communication processes between users and the system. We use different label sets ($S_1, S_2, \dots$) to represent the different concepts to be assessed in its filtering activity. These label sets S_i are chosen from those label sets that compose a LH , i.e., $S_i \in LH$. We should point out that the number of different label sets that we can use is limited by the number of levels of LH and therefore, in many cases the label sets S_i and S_j can be associated to a same label set of LH but with different interpretations depending on the concept to be modeled. We consider four concepts that can be assessed in the activity of the recommender system:

- *Importance degree* of a discipline with respect to a resource scope or user interest topic, which is assessed in S_1 .
- *Similarity degree* among resources or among users, which is assessed in S_2 .
- *Relevance degree* of a resource for a user, which is assessed in S_3 .
- *Satisfaction degree* expressed by a user to evaluate a recommended resource, which is assessed in S_4 .

Following the linguistic hierarchy shown in Fig. 2, we use the level 2 (5 labels) to represent importance degrees ($S_1 = S^5$), and the level 3 (9 labels) to represent similarity degrees ($S_2 = S^9$), relevance degrees ($S_3 = S^9$) and satisfaction degrees ($S_4 = S^9$). As the importance degrees are provided by TTO staff, we use a set of five labels to facilitate them the characterization of resource scopes or user interest topics. On the other hand, as the similarity and relevance degrees are computed automatically by the system we use the set of nine labels which presents an adequate granularity level to represent the results. Similarly, to provide users with a label set with an adequate granularity level we use the set of nine labels to express the satisfaction degrees. Using this LH , the linguistic terms in each level are the following:

- $S^5 = \{b_0 = \text{None} = N, b_1 = \text{Low} = L, b_2 = \text{Medium} = M, b_3 = \text{High} = H, b_4 = \text{Total} = T\}$
- $S^9 = \{c_0 = \text{None} = N, c_1 = \text{Very_Low} = VL, c_2 = \text{Low} = L, c_3 = \text{More_Less_Low} = MLL, c_4 = \text{Medium} = M, c_5 = \text{More_Less_High} = MLH, c_6 = \text{High} = H, c_7 = \text{Very_High} = VH, c_8 = \text{Total} = T\}$

In Fig. 3 we show the basic operating scheme of the recommender system which is based on four main components:

1. *Resources representation*. The system obtains an internal representation of the resources based on their scopes.
2. *User profiles representation*. The system obtains an internal representation of the user based on their research area and topics of interest.
3. *Recommendation process*. The system generates the recommendations according to the hybrid filtering approach.
4. *Feedback phase*. The users provides the system their opinions about the received recommendations.

In the following subsections we explain them in detail.

4.1. Resources representation

The resources we consider in our system are the research resources such as R&D bulletins, R&D&I, calls, notices or research projects. Once the TTO staff inserts all the available information about a new resource, the system obtains an internal representation mainly based on the resource scope. We use the *vector model* [44] to represent the resource scope and a classification composed by 25 disciplines (see Fig. 4), i.e., a research resource i is represented as

$$VR_i = (VR_{i1}, VR_{i2}, \dots, VR_{i25}),$$

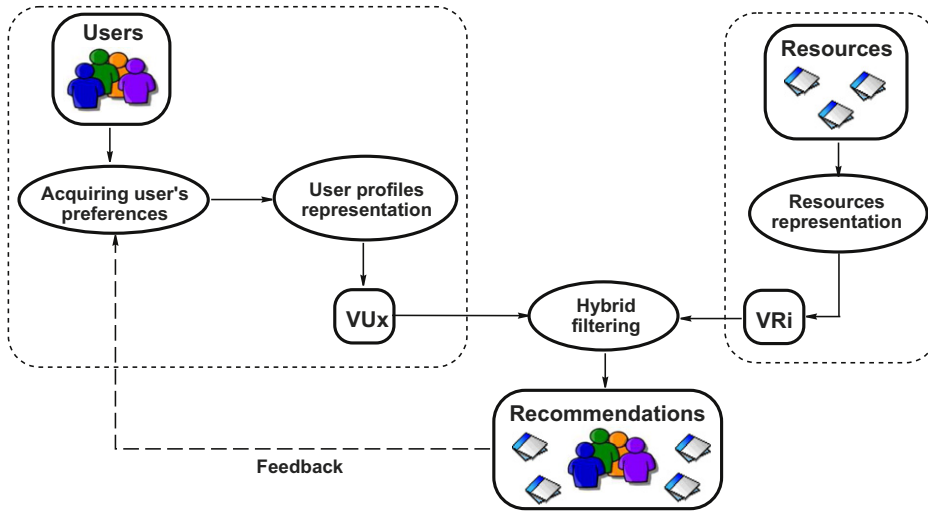


Fig. 3. Basic operating scheme.

☐ 1. Agricultural, animal breeding and fishing	☐ 2. Vegetal and animal biology and ecology
☐ 3. Biotechnology, molecular and cellular biology and genetics	☐ 4. Food science
☐ 5. Sscience and technology of materials	☐ 6. Earth science
☐ 7. Social science	☐ 8. Science and techonology of computers Total ▼
☐ 9. Laws	☐ 10. Economy
☐ 11. Energy and combustibles	☐ 12. Pharmacology and pharmacy
☐ 13. Philology and philosophy	☐ 14. Physics and space sciences
☐ 15. History and art	☐ 16. Civil engineering, transportation, construction and architecture
☐ 17. Industry, mechanics, naval and aeronautical engineering	☐ 18. Mathematics
☐ 19. Medicine and veterinary	☐ 20. Environment and environmental technology
☐ 21. Multidisciplinary	☐ 22. Scientific policy
☐ 23. Psychology and education sciences	☐ 24. Chemistry and chemical technology
☐ 25. Telecommunications, electrical, engineering, electronics and automatics Total ▼	

Total
 Total
 High
 Medium
 Low
 None

Fig. 4. Interface to define the disciplines of the resource scope or user preferences.

where each component $VR_{ij} \in S_1$ is a linguistic assessment that represents the importance degree of the discipline j with regard to the scope of i . These importance degrees are assigned by the TTO staff when they add new resources.

Example. We suppose that the TTO experts receive information about a call for an international conference on Computer Science. Then, one of the experts accesses to the application and inserts the new resource. He/she fills all the fields of the form and he/she uses the interface shown in Fig. 4 to select the disciplines of the resource scope. We assume that the expert selects the discipline titled “Science and technology of computers” with an importance degree “Total” and the discipline titled “Telecommunications, electrical engineering, electronics and automatics” with an importance degree “High”. These disciplines are in the positions 8 and 25 respectively of the used classification. The rest of disciplines have an importance degree with a value of “None”. So, the scope of the new resource is represented in the following way:

$$VR_i = (VR_{i1}, VR_{i2}, \dots, VR_{i25}),$$

where $VR_{i8} = (b_4, 0)$, $VR_{i25} = (b_3, 0)$ and the rest of positions have the value $(b_0, 0)$.

4.2. User profiles representation

We consider that the users of our system are the researchers of the university and the environment companies. To characterize a researcher the system stores the personal information (login, password, name, phone, email, etc.), research group (it is a string composed by 6 digits, 3 characters indicating the research area and 3 numbers identifying the group) and his/her topics of interest. Similarly, we use the vector model [44] to represent the topics of interest. Then, for a researcher e , we have a vector:

$$VU_e = (VU_{e1}, VU_{e2}, \dots, VU_{e25}),$$

where each component $VU_{ej} \in S_1$ is a linguistic assessment that represents the importance degree of the discipline j in the topics of interest of researcher e . Similarly these importance degrees are assigned by the TTO staff when they add a new researcher.

Example. We suppose that a user e has completed the registration form into the system. The user has inserted his/her personal information, research group, and his/her topics of interest. With this information a TTO expert inserts the user information into the system, and an internal representation of the user is obtained. For example, if e is a researcher in “economics” the vector representing his/her topics of interest would be the following:

$$VU_e = (VU_{e1}, VU_{e2}, \dots, VR_{e25}),$$

where $VU_{e10} = (b_4, 0)$, because the discipline titled “Economy” is in position 10 of the used classification. Similarly, the rest of positions have the value $(b_0, 0)$, because they have an importance degree of “None” for e .

Furthermore, to avoid the cold-starting problem to handle new items or new users [8,43], as in other systems (for example in Movielens [52]), when a new user is inserted, to confirm his/her register it is necessary that he/she assesses some of the resources stored in the system. To do this, the system is showing the items randomly and the user assesses what he/she wants.

4.3. Recommendation strategy

In this phase the system filters the incoming information to deliver it to the fitting users. This process is based on a matching process developed by similarity measures, such as Euclidean Distance or Cosine Measure [44]. In particular, we use the standard cosine measure but defined in a linguistic framework:

$$\sigma_l(V_1, V_2) = \Delta \left(g \times \frac{\sum_{k=1}^n (\Delta^{-1}(v_{1k}, \alpha_{v1k}) \times \Delta^{-1}(v_{2k}, \alpha_{v2k}))}{\sqrt{\sum_{k=1}^n (\Delta^{-1}(v_{1k}, \alpha_{v1k}))^2} \times \sqrt{\sum_{k=1}^n (\Delta^{-1}(v_{2k}, \alpha_{v2k}))^2}} \right)$$

with $\sigma_l(V_1, V_2) \in S_2 \times [-0.5, 0.5]$, and where g is the granularity of the term set used to express the relevance degree, i.e. S_2 , n is the number of disciplines and (v_{ik}, α_{vik}) is the 2-tuple linguistic value of discipline k in the vector V_i representing the resource scope or user interest topics, depending of the used filtering strategy.

This recommender system works according to a hybrid recommendation strategy. Burke [7] proposes a classification composed by different strategies according to the method of combining any two (or more) pure techniques to build a hybrid recommender system. In this sense, our proposal is based in a *switching hybrid approach*, which uses one technique or another, depending on some criterion. A system may try one technique and if the confidence of the results is not satisfactory, it may switch to another technique. In our system, depending on the case, a content-based recommendation approach or a collaborative one is applied. The former is applied when a new item is inserted and the latter is applied when a new researcher is registered. In both cases, the recommender system could send three types of recommendations to a researcher:

1. Research resources of his/her same area, i.e., specialized research resources.
2. Research resources of complementary areas, i.e. complementary research resources.
3. Collaboration possibilities with other researchers.

In the following, we explain both recommendation strategies.

4.3.1. Content-based recommendations

When a new resource i arrives to the system, the system calculates content-based recommendations to be sent to a researcher e as follows:

1. Compute the linguistic similarity degree between VR_i and VU_e .
2. Establish if the resource i could contribute to specialize or complement the researcher's profile. Assuming that $S_2 = S^9$, we consider that a resource i is related with the researcher's profile e if $\sigma_l(VR_i, VU_e) > (s_4^9, 0)$, i.e., if the linguistic similarity

degree is higher than the mid linguistic label. We consider that the resource i could contribute to specialize the researcher's profile e when $\sigma_l(VR_i, VU_e) \geq (s_6^9, 0)$. On the other hand, we consider that the resource i could contribute to complement the researcher's profile e when $(s_2^9, 0) \leq \sigma_l(VR_i, VU_e) < (s_6^9, 0)$.

3. If i is considered a specialization resource for e , then the system recommends this resource i to e with a relevance degree $i(e) \in S_3 \times [-0.5, 0.5]$ which is obtained as follows:

- (a) Look for all specialized research resources stored in the system that were previously assessed by e , i.e., the set of resources $K = \{1, \dots, k\}$ such that there exists the linguistic satisfaction assessment $e(j), j \in K$ and $\sigma_l(VU_e, VR_j) \geq (s_6^9, 0)$.
- (b) Then,

$$i(e) = \bar{x}_l^w(((e(1), 0), \sigma_l(VR_i, VR_1)), \dots, ((e(k), 0), \sigma_l(VR_i, VR_k))),$$

where $\bar{x}_l^w$ is the linguistic weighted average operator (see Definition 5).

4. If i is considered a complementary resource for e then the system recommends this resource i and its authors (community members that could be potential collaborators) to e with a relevance degree $i(e) \in S_3 \times [-0.5, 0.5]$ which is obtained as follows:

- (a) Look for all complementary research resources stored in the system that previously were assessed by e , i.e., the set of resources $K = \{1, \dots, k\}$ such that there exists the linguistic satisfaction assessment $e(j), j \in K$ and $(s_6^9, 0) > \sigma_l(VU_e, VR_j) \geq (s_2^9, 0)$. The latter defines a complementary linguistic interval around mid label that is considered the maximum complementary level.
- (b) Then,

$$i(e) = \bar{x}_l^w(((e(1), 0), f(i, 1)), \dots, ((e(k), 0), f(i, k))),$$

where f is a triangular multidisciplinary matching function that measures the complementary degree between two resources i and j ,

$$f(i, j) = \begin{cases} \Delta(2 \times \Delta^{-1}(\sigma_l(VR_i, VR_j))) & \text{if } 0 \leq \Delta^{-1}(\sigma_l(VR_i, VR_j)) \leq 1/2, \\ \Delta(2 \times (1 - \Delta^{-1}(\sigma_l(VR_i, VR_j)))) & \text{if } 1/2 < \Delta^{-1}(\sigma_l(VR_i, VR_j)) \leq 1. \end{cases}$$

4.3.2. Collaborative recommendations

When new users are inserted into the system, they receive recommendations about already inserted resources which may be interesting for them. Usually, new users provide little information about the items that satisfy their topics of interest, so we use the collaborative approach to generate their recommendations. Exactly, we follow a memory-based algorithm or nearest-neighbor algorithm, which generates the recommendations according to the preferences of nearest neighbors. This algorithm has proven good performance [22,70]. In the following we describe the process in detail.

Given a new researcher e , the recommendations to be sent to e are obtained in the following steps:

1. Identify the set of users $\aleph_e$ most similar to that new user e . To do so, we calculate the linguistic similarity degree between the topics of interest vector of the new user (VU_e) against the vectors of all users already inserted into the system ($VU_y, y = 1 \dots n$ where n is the number of users), that is, we calculate $\sigma_l(VU_e, VU_y) \in S_2$. As $S_2 = S^9$, we consider that the user y is near neighbor to e if $\sigma_l(VU_e, VU_y) > (s_6^9, 0)$, i.e., if the linguistic similarity degree is higher than the mid linguistic label.
2. Look for the resources stored in the system that were previously well assessed by the near neighbors of e , i.e., the set of resources $K = \{1, \dots, k\}$ such that there exists a linguistic satisfaction assessment $y(j), y \in \aleph_{e,j \in K}$, and $y(j) \geq (s_6^9, 0)$.
3. Discover if those resources could contribute with specialized or complementary formation. A resource $j \in K$ could contribute to specialize the researcher's formation e when $\sigma_l(VR_j, VU_e) \geq (s_6^9, 0)$. On the other hand, we consider that the resource j could contribute to complement the researcher's formation e when $(s_2^9, 0) \leq \sigma_l(VR_j, VU_e) < (s_6^9, 0)$.
4. If j is considered as a specialization resource for e , then the system recommends this resource j to e with a relevance degree $j(e) \in S_3 \times [-0.5, 0.5]$ which is obtained as follows:
 - (a) To look for all linguistic satisfaction assessments about resources that were well assessed by the nearest neighbors of e . That is, we recovery $y(j)$ with $j \in K$ and $y \in \aleph_e$.
 - (b) Then,

$$j(e) = \bar{x}_l^w(((y_1(j), 0), \sigma_l(VU_e, VU_{y_1})), \dots, ((y_n(j), 0), \sigma_l(VU_e, VU_{y_n}))),$$

where $y_1, \dots, y_n \in \aleph_e$ and $\bar{x}_l^w$ is the linguistic weighted average operator (see Definition 5).

5. If j is considered a complementary resource for e then the system recommends this resource j and its authors (community members that could be potential collaborators) to e with a relevance degree $j(e) \in S_3 \times [-0.5, 0.5]$ which is obtained as follows:

- (a) Look for all complementary research resources stored in the system that previously were well assessed by the nearest neighbors of e , i.e., the set of resources $K = \{1, \dots, k\}$ such that there exists the linguistic satisfaction assessment $y(j)$, with $j \in K, y \in \aleph_e$ and $(s_6^9, 0) > \sigma_l(VU_y, VR_j) \geq (s_2^9, 0)$. The latter defines a complementary linguistic interval around mid label that is considered the maximum complementary level.
- (b) Then,

$$j(e) = \bar{x}_l^w(((y_1(j), 0), h(e, y_1)), \dots, ((y_n(j), 0), h(e, y_n))),$$

where h is a triangular multidisciplinary matching function that measures the complementary degree between two researchers i and j ,

$$h(i,j) = \begin{cases} \Delta(2 \times \Delta^{-1}(\sigma_l(VR_i, VR_j))) & \text{if } 0 \leq \Delta^{-1}(\sigma_l(VU_i, VU_j)) \leq 1/2, \\ \Delta(2 \times (1 - \Delta^{-1}(\sigma_l(VR_i, VR_j))) & \text{if } 1/2 < \Delta^{-1}(\sigma_l(VU_i, VU_j)) \leq 1. \end{cases}$$

4.4. Feedback phase

In this phase the recommender system recalculates and updates the recommendations of the accessed resources. When the system sends recommendations to the users, then they provide a feedback by assessing the relevance of the recommendations, i.e., they supply their opinions about the recommendations received from the system. If they are satisfied with the received recommendation, they shall provide high values and viceversa. This feedback activity is developed in the following steps:

1. The system recommends the user U a resource R , and then the system asks him/her his/her opinion or evaluation judgements about recommended resource.
2. The user communicates his/her linguistic evaluation judgements to the system, $rc_y \in S_2$.
3. This evaluation is registered in the system for future recommendations. The system recalculates the linguistic recommendation of R by aggregating the opinions about R provided by all users. In such a way, the opinion supplied by U is considered. This can be done using the 2-tuple aggregation operator as $\bar{x}^e$ given in Definition 3.

5. Experiments and evaluation

In this section we present the evaluation of the proposed recommender system. We propose two kind of experiments, offline and online ones. We begin with an offline setting, where the proposed recommendation approach is compared with other approaches without user interaction, using a standard data set. However, in many applications, accurate predictions are important but insufficient with respect to the user satisfaction. For instance, users may be interested in discovering new items not expected for them, more than getting an exact prediction of their preferences. Consequently, we also propose online experiments, that is, practical studies where a small group of users interact with the system and report us their experiences.

5.1. Evaluation metrics

In the scope of recommender systems, precision, recall and F1 are widely used measures to evaluate the quality of the recommendations [11,15,64]. To calculate these metrics we need to build a contingency table to categorize the items with respect to the information needs (see Table 2). The items are classified both as relevant or irrelevant and selected (recommended to the user) or not selected.

Precision is defined as the ratio of the selected relevant items to the selected items, that is, it measures the probability of a selected item to be relevant:

$$P = \frac{N_{rs}}{N_s}.$$

Recall is calculated as the ratio of the selected relevant items to the relevant items, that is, it represents the probability of a relevant item to be selected:

$$R = \frac{N_{rs}}{N_r}.$$

F1 is a combination metric that gives equal weight to both precision and recall, and it is calculated as follows [11,64]:

$$F1 = \frac{2 \times R \times P}{R + P}.$$

Besides, in order to test the performance of our model and to compare it with other approaches, we also calculate the system accuracy, that is, its capability to predict users' ratings. We propose to use the **Mean Absolute Error (MAE)** [22,68], a

Table 2
Contingency table.

	Selected	Not selected	Total
Relevant	Nrs	Nrn	Nr
Irrelevant	Nis	Nin	Ni
Total	Ns	Nn	N

commonly used accuracy metric which considers the average absolute deviation between a predicted rating and the user's true rating:

$$MAE = \frac{\sum_{i=1}^n abs(p_i - r_i)}{n},$$

where n is the number of cases in the test set, p_i the predicted rating for a item, and r_i the true rating.

5.2. Offline experiments

In this subsection we present the offline experiments developed to analyze our system.

5.2.1. Data set

We have decided to use MovieLens data sets [21] to develop the offline experiments. We choose this option because the data sets are publicly available and have been usually used to evaluate recommender systems, and in such a way, we could compare our system with other models. MovieLens data sets are related with a cinematographic scope and they were collected by the GroupLens Research Project at the University of Minnesota during the seven-month period from September 19th, 1997 through April 22nd, 1998.

Specifically, we use the 100 K ratings data set which contains 1682 movies, 943 users and a total of 100000 ratings on a scale of 1–5 (where 1 = Awful, 2 = Fairly bad, 3 = Its OK, 4 = Will enjoy, 5 = Must see). Each user has rated at least 20 movies.

However, to apply this data set to our hybrid recommender system, we need to develop a transformation process in order to adapt the data to the features of our approach. In our system we represent both the resources and the user profiles using vectors. So, we need to transform the MovieLens data sets to this representation avoiding the loss of information. Then, we have to build vectors to represent the users' topics of interest and the movies. The idea is to obtain such vectors from the data stored in the MovieLens data sets.

The 1682 movies are classified into the following 19 genres: unknown, action, adventure, animation, children, comedy, crime, documentary, drama, fantasy, film-noir, horror, musical, mystery, romance, sci-fi, thriller, war and western. In fact, the file *u.item* contains information about the movies, with a tab separated list of the fields movie id, movie title, release date, video release date, IMDb URL, and the last 19 fields are the genres: a value of 1 indicates the movie is of that genre and a value of 0 indicates it is not; movies can be in several genres at once. For each movie we build a vector with 19 positions (one for each genre), following the approach pointed in subSection 4.1:

$$VR_i = (VR_{i1}, VR_{i2}, \dots, VR_{i19}),$$

where each component $VR_{ij} \in S_1$ is a linguistic assessment that represents the importance degree of the genre j with regard to the movie i . Therefore, when the value in the file *u.item* is 1 (the movie is of that genre), we assign the maximum label of S_1 ($(b_4, 0)$ in this case) and when the value is 0 the assigned label is the minimum of S_1 ($(b_0, 0)$).

On the other hand, our system works with the user topics of interest, which are also represented by a vector. So, for each user we need a vector similar to that used to represent the movies. The problem is that MovieLens data sets do not include this information directly, because the file *u.user* only includes demographic information about the users (user id, age, gender, occupation and zip code). However, the information about the topics of interest for each user could be obtained from the available data, aggregating the ratings assigned by the users on each movie with the genre information of the movies. The file *u.data* contains the 100000 ratings on a scale of 1–5; this is a tab separated list of user id, item id, rating and time-stamp. The information about the genres is in the file *u.item*; the movie ids are the ones used in the *u.data* data set. Following the approach pointed in subSection 4.2, for each user e we build a vector with 19 positions:

$$VU_e = (VU_{e1}, VU_{e2}, \dots, VU_{e19}),$$

where each component $VU_{ej} \in S_1$ is a linguistic assessment that represents the importance degree of the genre j in the topics of interest related with the user e . These importance degrees are calculated using a weighted average operator:

$$VU_{ej} = \Delta \left(\frac{\sum_{m=1}^n r_{em} \cdot g_{mj}}{\sum_{m=1}^n r_{em}} \right),$$

where r_{em} is the rating assigned by the user e on the movie m and g_{mj} is the value of the genre j for the movie m .

5.2.2. Results of offline experiments

We use the *cross validation* to determine the validity of our model and analyze the obtained results.

Cross validation is typically used to estimate how accurately a predictive model will perform in practice [62]. The data set is divided in complementary subsets, performing the analysis on one subset, called the training set, and validating the analysis on the other subset, called the testing set. To reduce variability, multiple rounds of cross-validation are performed using different partitions, and the validation results are averaged over the different rounds. In *k-fold cross validation* [62], the original sample is randomly partitioned into k folds. One fold is selected as the testing set, used to estimate the error, and the remaining $k - 1$ folds are used as training data set. The cross-validation process is then repeated k times, with each of the

folds used exactly once as the testing set. The k results then can be averaged to produce a single estimation about the deviations between the predictions and the actual ratings.

Values of the folding parameter k commonly assumed are 4, 5, ..., 10. We have chosen a value of $k = 5$. In order to perform 5-fold cross validation, we use the data sets $u1.base$ and $u1.test$ through $u5.base$ and $u5.test$ provided by MovieLens which split the collection into 80% for training and 20% for testing, respectively. From the training data sets we build the necessary vectors as we have shown in the previous section. We use them as the input data to predict the unrated ratings. It allows us to measure the system capability in order to predict the users' ratings, calculating the MAE. Besides, to test the effect of the number of neighbors (value of N_e used in the collaborative recommendations) on the accuracy of the system, we have considered the most 5, 10, 20, 30 and 50 similar users. The obtained results are shown in the Table 3 and illustrated in Fig. 5.

As we can see, the performance of the system is quite uniform across the MovieLens data set, but considering 10 and 20 similar users we obtain a better average MAE than the rest of configurations. The other three groups (with the most 5, 30 and 50 similar users), present results close to one another. Fig. 5 clearly shows that the average MAE increases as the number of neighbors grows or when we consider very few neighbors. When the number of neighbors is between 10 and 20, there is a significant drop in average MAE which indicates a considerable increase in prediction accuracy; in fact, the best average results are obtained considering the most 20 similar users. Therefore, we decide that a number of neighbors between 15 and 20, are the most suitable for our system.

5.2.3. Comparison with other approaches

In order to compare the results of our system with other, we have implemented several content-based and collaborative models. Firstly, we have implemented a pure content-based approach (CB) [4,5,57] in which the similarity between two items is calculated using the cosine measure. We also have implemented the user-based collaborative approach (UBC) [21,64,70]. This method uses the ratings of users that are most similar to the target user for predicting the ratings of unrated items; the similarity between users is computed using Pearson's correlation coefficient. Finally, we have implemented the item-based collaborative approach (IBC) [4,17,65] in which the similarities of items are used to predict the ratings. The prediction is computed by taking a weighted average of the target user's ratings on similar items. In our experimentation we have used both the cosine and Pearson measure, titled IBC-C and IBC-P, respectively.

To compare the different approaches, we have followed the experimental setting described previously, that is, we perform the 5-fold cross validation, using as training and testing data sets the files $u1.base$ and $u1.test$ through $u5.base$ and $u5.test$ provided by MovieLens. With these experiments we calculate the average MAE for all the tests and rounds. To do the comparison, in the case of our system we have used the average MAE for the five values of N_e studied in the previous subsection (see Table 3). We prefer to use the average value and not the better MAE, to obtain more significant and realistic results. Table 4 presents the MAE results obtained by each approach, where we can see how our system improves the results

Table 3

MAE values for our system with MoveLens data sets.

N_e	u1	u2	u3	u4	u5	Average MAE
5	0.7405	0.7398	0.7424	0.7432	0.7437	0.7419
10	0.7379	0.7339	0.7353	0.7386	0.7381	0.7368
20	0.7356	0.7351	0.7357	0.7372	0.7370	0.7361
30	0.7434	0.7431	0.7446	0.7452	0.7448	0.7442
50	0.7471	0.7463	0.7468	0.7479	0.7473	0.7471

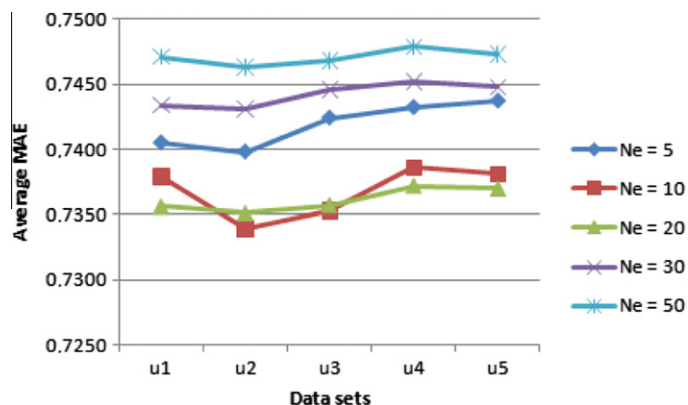


Fig. 5. MAE values for our system.

Table 4

Average MAE values to compare with other models.

	Our system	CB	UBC	IBC-C	IBC-P
Average MAE	0.7412	0.9187	0.7848	0.7705	0.7716
Improvement %		23.94%	5.88%	3.95%	4.10%

obtained by the rest of approaches. The row entitled with *Improvement %* presents the improvement percentage obtained with our system over the other approaches.

5.3. Online experiments

In this subsection we present the online experiments followed to analyze the system performance. We have enabled both the new recommender system and SIRE2IN for a small group of users, who interact with the system and report us their experience.

5.3.1. Data set

For the online evaluation, we have considered a data set with 200 research resources related with different areas collected by the TTO staff from different information sources. These resources were included into the system following the indications described in Section 4.3.1. We assume that the recommender system has to generate recommendations to 15 users and that these users have completed the registration process and evaluated at least 25 resources. From these user assessments, the system is able to build the user profiles.

The resources and the provided user assessments constitute our training data set. Then, we have added 100 new research resources that conform the test data set. The system filtered these 100 resources and it recommended them to the suitable users. To obtain data to compare, these 100 new research resources also were recommended using the advices of the TTO staff.

5.3.2. A comparative study

In a first experiment we present a comparative study between our new recommender system and SIRE2IN. We have decided to use SIRE2IN to do this comparison because both systems work in the same framework, and therefore, the experimental setting is the same. In this experiment we only compare the recommendation approaches with the following restrictions: we do not consider the collaboration recommendation possibilities of our new recommender system neither distinctions between specialized and complementary resources (we only take into account if a resource is recommended or not).

By comparing the recommendations generated by our new recommender system with those generated by the TTO staff we obtained the contingency data given in Table 5. For example, for user 1, the new recommender system selected 17 research resources as relevant. However, from the information provided by the TTO staff, we can see that the system selected 9 irrelevant resources for user 1, and it was not able to select 6 resources that TTO staff considered relevant for the user. In a similar way, we obtained the contingency table for SIRE2IN.

In Tables 6 and 7 we show the evaluation metrics for SIRE2IN and our new recommender system, respectively. As we can see, with our hybrid recommender system the indicator F1 was 68.11%, which was greater than that obtained for SIRE2IN (54.49%). Therefore, the new system worked better than SIRE2IN. Fig. 6 shows a graph with the F1 value for both systems.

Table 5

Experimental contingency table for the new system.

	Nrs	Nrn	Nis	Nr	Ns
User1	17	6	9	23	26
User2	15	8	9	23	24
User3	18	10	10	28	28
User4	14	5	6	19	20
User5	22	11	8	33	30
User6	26	9	14	35	40
User7	19	8	7	27	26
User8	17	9	6	26	23
User9	22	10	14	32	36
User10	23	9	12	32	35
User11	20	11	10	31	30
User12	24	9	11	33	35
User13	18	8	10	26	28
User14	17	9	7	26	24
User15	18	8	9	26	27

Table 6
Metrics for SIRE2IN.

	Precision (%)	Recall (%)	F1 (%)
User1	62.50	60.00	61.22
User2	70.00	56.00	62.22
User3	45.00	60.00	51.43
User4	41.67	52.63	46.51
User5	50.00	40.00	44.44
User6	54.55	60.00	57.14
User7	66.67	52.63	58.82
User8	55.00	52.38	53.66
User9	54.55	54.55	54.55
User10	54.55	52.17	53.33
User11	58.33	53.85	56.00
User12	60.00	51.72	55.56
User13	57.14	52.17	54.55
User14	56.25	52.94	54.55
User15	54.55	52.17	53.33
Average	56.05	53.55	54.49

Table 7
Metrics for the new system.

	Precision (%)	Recall (%)	F1 (%)
User1	65.38	73.91	69.39
User2	62.50	65.22	63.83
User3	64.29	64.29	64.29
User4	70.00	73.68	71.79
User5	73.33	66.67	69.84
User6	65.00	74.29	69.33
User7	73.08	70.37	71.70
User8	73.91	65.38	69.39
User9	61.11	68.75	64.71
User10	65.71	71.88	68.66
User11	66.67	64.52	65.57
User12	68.57	72.73	70.59
User13	64.29	69.23	66.67
User14	70.83	65.38	68.00
User15	66.67	69.23	67.92
Average	67.42	69.03	68.11

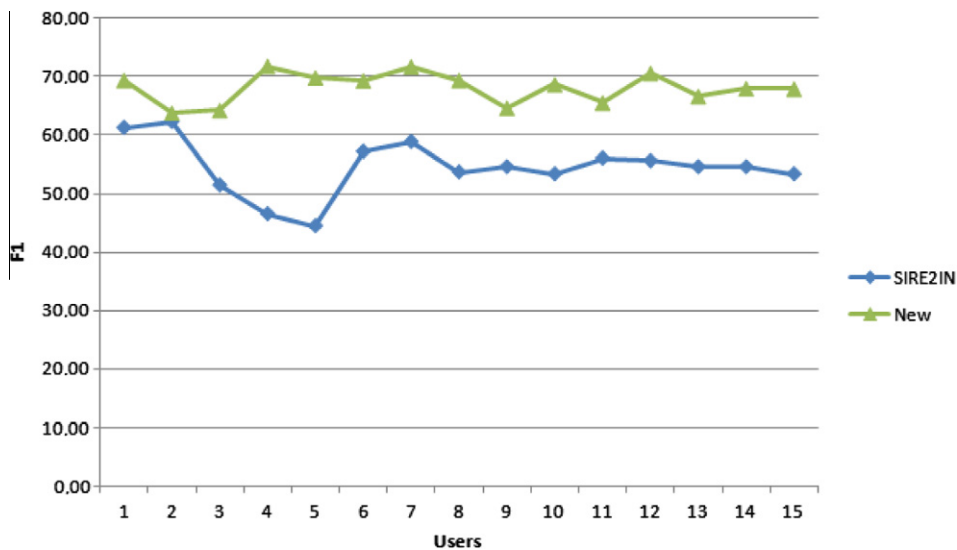


Fig. 6. Comparative graph of performance for SIRE2IN and the new system.

5.3.3. Evaluating the new functionalities of the recommender system

To complete the evaluation of our new recommender system, we test its new functionalities, i.e., its capacities to discover both specialized or complementary resources and collaboration possibilities.

We used the online experiment data set outlined previously, and thus, the training data set was composed by 200 research resources of different areas and 15 user profiles. Then, we added 100 new resources that constituted the test data set. The system filtered the 100 resources and recommended them to the suitable users, but now indicating if the resource was specialized or complementary. These 100 resources were also recommended and classified as specialized or complementary by the TTO staff. Using these data we obtained the contingency table shown in Table 8. For example, for user 1, the TTO staff considered 23 relevant resources, of which 16 were specialized and 7 were complementary. Our system selected 26 resources as relevant for user 1, being only 17 really relevant. From these 17 relevant resources, the system classified 15 as specialized and 2 as complementary. Comparing with the recommendations provided by the TTO staff we had 2 resources which are misclassified. So, the success rate for the user 1 was $((17 - 2)/26) \times 100 = 57.69\%$. Analyzing the data given in Table 8, we detected that the system shows an average precision (success rate) of 61.28%.

If we focus on the value of precision, we can see that this value has declined in comparison with the precision obtained by the system when its discovering possibilities are not considered. In Table 7 we observe that the system obtains a precision value of 67.42%. However, this value of precision is greater than 56.05% obtained for SIRE2IN (see Table 6).

Similarly, we used the previous scenario to analyze the collaboration possibilities of our new recommender system. However, in this case, the items to recommend are not the research resources, but the collaboration opportunities that could appear when the resource is a research project. Thus, we assumed that our system had to recommend research resources to 15 users and a training data set composed by 200 research resources of different areas. Then, we added 100 new resources, of which 30 resources were research projects that constituted the test data set. To compare the collaboration recommendations provided by the system and by the TTO staff we used those 30 projects, and not only the projects considered as relevant by the system or by the TTO staff. So, we can obtain specific measures with regard to collaboration recommendations.

Then, for the 30 projects we compared the collaboration recommendations made by the system with the collaboration recommendations provided by the TTO staff. We classified the collaboration recommendations taking into account the categorization described in Table 9. To understand the meaning of this table we provide the following example. Suppose that for project 1, the TTO staff selected user 7 and indicated him/her that he/she could collaborate with users 2, 11 and 12 to develop the project. Our system also selected user 7 for project 1, but in this case it recommended the collaboration with users 2, 3 and 12, and therefore, these recommendations did not match with the TTO staff recommendations. That is, our system presented 2 hits (for users 2 and 12), 1 failure (user 3) and a non-detected collaboration (user 11). Then, for project 1, $Nchs = 2$, $Nchn = 1$ and $Ncfs = 1$. Assuming this framework, we obtained the Table 10 for the 30 projects, being the average precision of 70.44%, the average recall of 72.50% and an average F1 of 70.40%, which show a satisfactory behavior of our system.

Table 8
Success rates in classifying.

	Success rate (%)
User1	57.69
User2	54.17
User3	53.57
User4	70.00
User5	60.00
User6	52.50
User7	65.38
User8	69.57
User9	55.56
User10	62.86
User11	63.33
User12	60.00
User13	60.71
User14	70.83
User15	62.96
Average	61.28

Table 9
Contingency table for the collaboration recommendations.

	Selected	Not selected	Total
Considered by TTO	Nchs	Nchn	Nch
Not considered by TTO	Ncfs	Ncfn	Ncf
Total	Ns	Nn	N

Table 10

Experimental contingency table for the collaboration recommendations.

Project	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
Nchs	2	3	4	2	4	2	3	2	3	3	2	2	3	2	2
Nchn	1	1	2	1	1	0	1	1	1	2	2	1	1	0	1
Ncfs	1	1	1	2	2	1	1	1	1	1	1	1	0	2	1
Nch	3	4	6	3	5	2	4	3	4	5	4	3	4	2	3
Ns	3	4	5	4	6	3	4	3	4	4	3	3	3	4	3
Project	16	17	18	19	20	21	22	23	24	25	26	27	28	29	30
Nchs	3	2	2	3	4	3	2	3	3	3	2	3	2	3	1
Nchn	2	1	1	1	1	1	1	2	1	2	0	1	0	1	1
Ncfs	1	1	0	1	1	1	1	1	2	1	1	1	1	2	1
Nch	5	3	3	4	5	4	3	5	4	5	2	4	2	4	2
Ns	4	3	2	4	5	4	3	4	5	4	3	4	3	5	2

In this case, we cannot compare these results with the obtained using SIRE2IN, because SIRE2IN looks for community members to collaborate in a very restricted way. SIRE2IN recommends researchers who present a similar profile to the user. However, the obtained results indicate that the collaboration recommendations provided by our system are useful to researchers, and quite similar to those provided by the TTO staff.

6. Concluding remarks

The TTO is responsible for putting into action and managing the activities which generate knowledge and technical and scientific collaboration. A service that is particularly important to fulfill this objective is the selective dissemination of information about research resources. The TTO staff and researchers need tools to assist them in their processes of information discovering because of the large amount of information available on these systems.

We have presented a new fuzzy linguistic recommender system to spread selectively research resources in a TTO, that solves the problems encountered in previous proposals. Particularly, we propose to replace the recommendation engine using a hybrid approach, that is, integrating a content-based approach with a collaborative one, in order to take the advantages of both strategies and reduce the disadvantages of each one of them. This new recommender system incorporates new functionalities, recommending specialized resources, complementary resources and collaboration possibilities that allows the researchers to meet other researchers and to form multidisciplinary groups. Besides, the system improves the feedback process using satisfaction degrees.

We have applied our research in a real environment provided by the TTO. The system advises researchers and environment companies about resources that could be interesting for them and collaboration possibilities with other researchers. The experimental results show us significant improvements over previous proposals.

Analyzing our system, we could conclude that its main limitation is the need for interaction with TTO staff to establish the internal representations for the user profiles and the items. With regard to future research, we believe that a promising direction is to study automatic techniques to establish the representation of user profiles and items. Moreover, we want to explore new improvements of the recommendation approach, exploring new methodologies for the generation of recommendations, as for example, bibliometric tools to enrich the information on the researchers and research resources [1,16].

Acknowledgments

This paper has been developed with the financing of Projects 90/07, TIN2007-61079, PET2007-0460, TIN2010-17876, TIC5299 and TIC-5991.

References

- [1] S. Alonso, F.J. Cabrerizo, E. Herrera-Viedma, F. Herrera, hg-index: a new index to characterize the scientific output of researchers based on the h-and g-Indices, *Scientometrics* 82 (2) (2010) 391–400.
- [2] S. Alonso, E. Herrera-Viedma, F. Chiclana, F. Herrera, A web based consensus support system for group decision making problems and incomplete preferences, *Information Sciences* 180 (23) (2010) 4477–4495.
- [3] B. Arfi, Fuzzy decision making in politics, *A Linguistic Fuzzy-set Approach (LFSA)*. *Political Analysis* 13 (1) (2005) 23–56.
- [4] A.B. Barragáns-Martínez, E. Costa-Montenegro, J.C. Burguillo, M. Rey-López, F.A. Mikic-Fonte, A. Peleteiro, A hybrid content-based and item-based collaborative filtering approach to recommend TV programs enhanced with singular value decomposition, *Information Sciences* 180 (22) (2010) 4290–4311.
- [5] C. Basu, H. Hirsh, W. Cohen, Recommendation as classification: using social and content-based information in recommendation, in: *Proceedings of the Fifteenth National Conference on Artificial Intelligence*, 1998, pp. 714–720.
- [6] G. Bordogna, G. Pasi, A fuzzy linguistic approach generalizing boolean information retrieval: a model and its evaluation, *Journal of the American Society for Information Science* 44 (1993) 70–82.
- [7] R. Burke, Hybrid recommender systems: Survey and Experiments, *User Modeling and User-Adapted Interaction* 12 (4) (2002) 331–370.

- [8] R. Burke, Hybrid web recommender systems, in: P. Brusilovsky, A. Kobsa, W. Nejdl (Eds.): The Adaptive Web, LNCS 4321, 2007, pp. 377–408.
- [9] F.J. Cabrerizo, J. López-Gijón, A.A. Ruiz-Rodríguez, E. Herrera-Viedma, A model based on fuzzy linguistic information to evaluate the quality of digital libraries, *International Journal of Information Technology and Decision Making* 9 (3) (2010) 455–472.
- [10] F.J. Cabrerizo, I.J. Pérez, E. Herrera-Viedma, Managing the Consensus in Group Decision Making in an Unbalanced Fuzzy Linguistic Context with Incomplete Information, *Knowledge-Based Systems* 23 (2) (2010) 169–181.
- [11] Y. Cao, Y. Li, An intelligent fuzzy-based recommendation system for consumer electronic products, *Expert Systems with Applications* 33 (2007) 230–240.
- [12] S.L. Chang, R.C. Wang, S.Y. Wang, Applying a direct multigranularity linguistic and strategy-oriented aggregation approach on the assessment of supply performance, *European Journal of Operational Research* 177 (2) (2007) 1013–1025.
- [13] Z. Chen, D. Ben-Arieh, On the fusion of multi-granularity linguistic label sets in group decision making, *Computers and Industrial Engineering* 51 (3) (2006) 526–541.
- [14] M. Claypool, A. Gokhale, T. Miranda, Combining content-based and collaborative filters in an online newspaper, in: *Proceedings of the ACM SIGIR-99 Workshop on Recommender Systems-Implementation and Evaluation*, 1999.
- [15] C.W. Cleverdon, E.M. Keen, Factors determining the performance of indexing systems, vol. 2-Test Results, in: *ASLIB Cranfield Res. Proj.*, Cranfield, Bedford, England, 1966.
- [16] M.J. Cobo, A.G. López-Herrera, E. Herrera-Viedma, F. Herrera, An Approach for Detecting, Quantifying, and Visualizing the Evolution of a Research Field: A Practical Application to the Fuzzy Sets Theory Field, *Journal of Informetrics* 5 (1) (2011) 146–166.
- [17] M. Deshpande, G. Karypis, Item-Based Top-N Recommendation algorithms, *ACM Transactions on Information Systems* 22 (1) (2004) 143–177.
- [18] L. Duen-Ren, L. Chin-Hui, L. Wang-Jung, A hybrid of sequential rules and collaborative filtering for product recommendation, *Information Sciences* 179 (2009) 3505–3519.
- [19] N. Good, J.B. Schafer, J.A. Konstan, A. Borchers, B.M. Sarwar, J.L. Herlocker, J. Riedl, Combining collaborative filtering with personal agents for better recommendations, in: *Proceedings of the Sixteenth National Conference on Artificial Intelligence*, 1999, pp. 439–446.
- [20] U. Hanani, B. Shapira, P. Shoval, Information filtering: overview of issues, research and systems, *User Modeling and User-Adapted Interaction* 11 (2001) 203–259.
- [21] J. Herlocker, J. Konstan, A. Borchers, J. Riedl, An algorithmic framework for performing collaborative filtering, in: *Proceedings of the 22nd Annual International ACM SIGIR Conference on Research and Development in Information Retrieval*, 1999, pp. 230–237.
- [22] J.L. Herlocker, J.A. Konstan, L.G. Terveen, J.T. Riedl, Evaluating collaborative filtering recommender systems, *ACM Transactions on Information Systems* 22 (1) (2004) 5–53.
- [23] F. Herrera, E. Herrera-Viedma, Aggregation operators for linguistic weighted information, *IEEE Transactions on Systems, Man and Cybernetics, Part A: Systems* 27 (1997) 646–656.
- [24] F. Herrera, E. Herrera-Viedma, Linguistic decision analysis: steps for solving decision problems under linguistic information, *Fuzzy Sets and Systems* 115 (2000) 67–82.
- [25] F. Herrera, E. Herrera-Viedma, L. Martínez, A fusion approach for managing multi-granularity linguistic term sets in decision making, *Fuzzy Sets and Systems* 114 (2000) 43–58.
- [26] F. Herrera, E. Herrera-Viedma, J.L. Verdegay, Direct approach processes in group decision making using linguistic OWA operators, *Fuzzy Sets and Systems* 79 (1996) 175–190.
- [27] F. Herrera, L. Martínez, A 2-tuple fuzzy linguistic representation model for computing with words, *IEEE Transactions on Fuzzy Systems* 8 (6) (2000) 746–752.
- [28] F. Herrera, L. Martínez, A model based on linguistic 2-tuples for dealing with multigranularity hierarchical linguistic contexts in multiexpert decision-making, *IEEE Transactions on Systems, Man and Cybernetics, Part B: Cybernetics* 31 (2) (2001) 227–234.
- [29] E. Herrera-Viedma, Modeling the retrieval process of an information retrieval system using an ordinal fuzzy linguistic approach, *Journal of the American Society for Information Science and Technology* 52 (6) (2001) 460–475.
- [30] E. Herrera-Viedma, An information retrieval system with ordinal linguistic weighted queries based on two weighting elements, *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems* 9 (2001) 77–88.
- [31] E. Herrera-Viedma, O. Cordon, M. Luque, A.G. López, A.M. Muñoz, A model of fuzzy linguistic IRS based on multi-granular linguistic information, *International Journal of Approximate Reasoning* 34 (3) (2003) 221–239.
- [32] E. Herrera-Viedma, A.G. López-Herrera, A model of information retrieval system with unbalanced fuzzy linguistic information, *International Journal of Intelligent Systems* 22 (11) (2007) 1197–1214.
- [33] E. Herrera-Viedma, A.G. López-Herrera, A review on information accessing systems based on fuzzy linguistic modelling, *International Journal of Computational Intelligence Systems* 3 (4) (2010) 420–437.
- [34] E. Herrera-Viedma, A.G. López-Herrera, S. Alonso, J.M. Moreno, F.J. Cabrerizo, C. Porcel, A computer-supported learning system to help teachers to teach fuzzy information retrieval systems, *Information Retrieval* 12 (2) (2009) 179–200.
- [35] E. Herrera-Viedma, A.G. López-Herrera, M. Luque, C. Porcel, A fuzzy linguistic IRS model based on a 2-tuple fuzzy linguistic approach, *International Journal of Uncertainty, Fuzziness and Knowledge-based Systems* 15 (2) (2007) 225–250.
- [36] E. Herrera-Viedma, A.G. López-Herrera, C. Porcel, Tuning the matching function for a threshold weighting semantics in a linguistic information retrieval system, *International Journal of Intelligent Systems* 20 (9) (2005) 921–937.
- [37] E. Herrera-Viedma, L. Martínez, F. Mata, F. Chiclana, A consensus support system model for group decision-making problems with multi-granular linguistic preference relations, *IEEE Transactions on Fuzzy Systems* 13 (5) (2005) 644–658.
- [38] E. Herrera-Viedma, G. Pasi, A.G. López-Herrera, C. Porcel, Evaluating the information quality of web sites: a methodology based on fuzzy computing with words, *Journal of the American Society for Information Science and Technology* 57 (4) (2006) 538–549.
- [39] E. Herrera-Viedma, E. Peis, Evaluating the informative quality of documents in SGML-format using fuzzy linguistic techniques based on computing with words, *Information Processing and Management* 39 (2) (2003) 233–249.
- [40] A. Iskold, The Art, Science and Business of Recommendation Engines, January 2007. <http://www.readwriteweb.com/archives/recommendation_engines.php>.
- [41] M.H. Hsu, A personalized English learning recommender system for ESL students, *Expert Systems with Applications* 34 (2008) 683–688.
- [42] Y. Jiang, Z. Fan, J. Ma, A method for group decision making with multi-granularity linguistic assessment information, *Information Sciences* 178 (4) (2008) 1098–1109.
- [43] H. Jun Ahn, A new similarity measure for collaborative filtering to alleviate the new user cold-starting problem, *Information Sciences* 178 (1) (2008) 37–51.
- [44] R.R. Korfhage, *Information Storage and Retrieval*, Wiley Computer Publishing, New York, 1997.
- [45] S.K. Lee, Y.H. Cho, S.H. Kim, Collaborative filtering with ordinal scale-based implicit ratings for mobile music recommendations, *Information Sciences* 180 (11) (2010) 2142–2155.
- [46] C. Long-Sheng, H. Fei-Hao, C. Mu-Chen, H. Yuan-Chia, Developing recommender systems with the consideration of product profitability for sellers, *Information Sciences* 178 (4) (2008) 1032–1048.
- [47] A.G. López-Herrera, E. Herrera-Viedma, F. Herrera, Applying multi-objective evolutionary algorithms to the automatic learning of extended boolean queries in fuzzy ordinal linguistic information retrieval systems, *Fuzzy Sets and Systems* 160 (2009) 2192–2205.
- [48] I. Macho-Stadler, D. Perez-Castrillo, R. Veuglers, Licensing of university inventions: the role of a technology transfer office, *International Journal of Industrial Organization* 25 (2007) 483–510.
- [49] G. Marchionini, Research and Development in Digital Libraries. <http://ils.unc.edu/march/digital_library_R_and_D.html>. Last access: August, 2011.

- [50] G.D. Markman, P.H. Phan, D.B. Balkin, P.T. Gianiodis, Entrepreneurship and university-based technology transfer, *Journal of Business Venturing* 20 (2005) 241–263.
- [51] F. Mata, L. Martínez, E. Herrera-Viedma, An adaptive consensus support model for group decision making problems in a multi-granular fuzzy linguistic context, *IEEE Transactions on Fuzzy Systems* 17 (2) (2009) 279–290.
- [52] MovieLens-movie recommendations. <<http://movielens.umn.edu/login>>.
- [53] J.M. Morales del Castillo, E. Peis, E. Herrera-Viedma, A filtering and recommender system for E-scholars, *International Journal of Technology of Enhanced Learning* 2 (3) (2010) 227–240.
- [54] J.M. Morales del Castillo, E. Peis, A.A. Ruiz-Rodríguez, E. Herrera-Viedma, Recommending biomedical resources: a fuzzy linguistic approach based on semantic web, *International Journal of Intelligent System* 25 (2010) 1143–1157.
- [55] J.M. Moreno, J.M. Morales del Castillo, C. Porcel, E. Herrera-Viedma, A quality evaluation methodology for health-related websites based on a 2-tuple fuzzy linguistic approach, *Soft Computing* 14 (8) (2010) 887–897.
- [56] E. Nasibov, A. Kinay, An iterative approach for estimation of student performances based on linguistic evaluations, *Information Sciences* 179 (5) (2009) 688–698.
- [57] A. Popescul, L.H. Ungar, D.M. Pennock S. Lawrence, Probabilistic models for unified collaborative and content-based recommendation in sparse-data environments, in: *Proceedings of the Seventeenth Conference on Uncertainty in Artificial Intelligence (UAI)*, San Francisco, 2001, pp. 437–444.
- [58] C. Porcel, E. Herrera-Viedma, Dealing with incomplete information in a fuzzy linguistic recommender system to disseminate information in university digital libraries, *Knowledge-Based Systems* 23 (2010) 32–39.
- [59] C. Porcel, A.G. López-Herrera, E. Herrera-Viedma, A recommender system for research resources based on fuzzy linguistic modeling, *Expert Systems with Applications* 36 (2008) 5173–5183.
- [60] C. Porcel, J.M. Morales del Castillo, M.J. Cobo, A.A. Ruiz-Rodríguez, E. Herrera-Viedma, An improved recommender system to avoid the persistent information overload in a university digital library, *Control and Cybernetics* 39 (4) (2010) 899–924.
- [61] L.M. Quiroga, J. Mostafa, An experiment in building profiles in information filtering: the role of context of user relevance feedback, *Information Processing and Management* 38 (2002) 671–694.
- [62] P. Refaailzadeh, L. Tang, H. Liu, Cross-Validation. <<http://www.public.asu.edu/ltang9/papers/ency-cross-validation.pdf>>.
- [63] P. Reisman, H.R. Varian, Recommender systems, *Special issue of Communications of the ACM* 40 (3) (1997) 56–59.
- [64] B. Sarwar, G. Karypis, J. Konstan, J. Riedl, Analysis of recommendation algorithms for e-commerce, in: *Proceedings of ACM E-Commerce 2000 conference*, 2000, pp. 158–167.
- [65] B. Sarwar, G. Karypis, J. Konstan, J. Riedl, Item-based collaborative filtering recommendation algorithms, in: *Proceedings of the ACM World Wide Web Conference*, 285295, 2001.
- [66] R. Schenkel, T. Crecelius, M. Kacimi, T. Neumann, J.X. Parreira, M. Spaniol, G. Gerhard Weikum, Social wisdom for search and recommendation, *Data Engineering Bulletin* 31 (2) (2008) 40–49.
- [67] J. Serrano-Guerrero, E. Herrera-Viedma, J.A. Olivas, A. Cerezo, F.P. Romero, A google wave-based fuzzy recommender system to disseminate information in university digital libraries 2.0, *Information Sciences* 181 (2011) 1503–1516.
- [68] G. Shani, A. Gunawardana, Evaluating Recommendation Systems, in: *Francesco Ricci, Lior Rokach, Bracha Shapira, Paul B. Kantor (Eds.), Recommender Systems Handbook*, Springer, 2011.
- [69] D.S. Siegel, D.A. Waldman, L.E. Atwater, A.N. Link, Toward a model of the effective transfer of scientific knowledge from academicians to practitioners: qualitative evidence from the commercialization of university technologies, *Journal of Engineering and Technology Management* 21 (2004) 115–142.
- [70] P. Symeonidis, A. Nanopoulos, A.N. Papadopoulos, Y. Manolopoulos, Collaborative recommender systems: combining effectiveness and efficiency, *Expert Systems with Applications* 34 (2008) 2995–3013.
- [71] I. Truck, H. Akdag, A tool for aggregation with words, *Information Sciences* 179 (14) (2009) 2317–2324.
- [72] R.R. Yager, Fuzzy logic methods in recommender systems, *Fuzzy Sets and Systems* 136 (2) (2003) 133–149.
- [73] L.A. Zadeh, The concept of a linguistic variable and its applications to approximate reasoning, Part I. *Information Sciences* 8 (1975) 199–249. Part II, *Information Sciences*, 8, 1975, pp. 301–357. Part III, *Information Sciences*, 9, 1975, pp. 43–80.
- [74] A. Zenebe, A.F. Nocio, Representation, similarity measures and aggregation methods using fuzzy sets for content-based recommender systems, *Fuzzy Sets and Systems* 160 (1) (2009) 76–94.
- [75] G. Zhang, J. Lu, A linguistic intelligent user guide for method selection in multi-objective decision support systems, *Information Sciences* 179 (14) (2009) 2299–2308.