

A Reviewer Recommendation System Based on Collaborative Intelligence

Kai-Hsiang Yang², Tai-Liang Kuo¹, Hahn-Ming Lee^{1,2}, Jan-Ming Ho²

¹Department of Computer Science and Information Engineering,
National Taiwan University of Science and Technology, Taipei, Taiwan
Email: {M9615006, hmlee}@mail.ntust.edu.tw

²Institute of Information Science, Academia Sinica, Taipei, Taiwan
Email: {khyang, hmlee, hoho}@iis.sinica.edu.tw

Abstract

In this paper, expert-finding problem is transformed to a classification issue. We build a knowledge database to represent the expertise characteristic of domain from web information constructed by collaborative intelligence, and an incremental learning method is proposed to update the database. Furthermore, results are ranked by measuring the correlation in the concept network from online encyclopedia. In our experiments, we use the real world dataset which comprise 2,701 experts who are categorized into 8 expertise domains. Our experimental results show that the expertise knowledge extracted from collaborative intelligence can improve efficiency and effect of classification and increase the precision of ranking expert at least 20%.

1 Introduction

The reviewer recommendation is an important but complex system [7]. The key problem of reviewer recommendation is to identify experts for specific topics [5, 10]. It considers an expert who had enough expertise for that specific topic. The problem of expert finding has been mentioned in previous work [1, 2, 5, 7, 10] and the experts are identified by expertise modeling from online communities [11] or their publication [9]. Statistics of keywords co-occurrence in documents or publications [1, 5] are approached to find the similar documents to queries, and the authors as the experts for queries. Ontology-based approach for expertise matching is more efficiently and effectively [3, 9]. But the main drawback of ontology-based approach is that it needs lots of effort to construct and maintain the ontology in many domains. However, the increased new terms followed the increase of the development of research field. Approach differed from keyword co-occurrence is consideration for degree of activity and the category and type of documents

in online community [2, 11]. We overcome these drawbacks by using the online encyclopedia as the semantic kernel [4, 6] to construct our Expertise Knowledge Database (*EKD*) by an incremental learning method. The online encyclopedia is named Wikipedia which is built by collaborative intelligence from all over the world. The *EKD* can help us modeling the characters of domains and classifying the proposal into related domains. The Wikipedia category network is used as the Wikipedia Concept Network (*WCN*) to compute the word-semantic relatedness.

In this paper, we propose a approach to solve these issue in a real world task which is a peer review process for review proposals. Peer review is an essential but tough task for research councils, journal editors, and conference program chairs [7]. Besides, many research proposals are multidisciplinary in the computer science domain. (e.g., some proposals address predictions of stock quote by rule based machine learning technologies.) It is a challenge to find suitable experts efficiently and it needs many information to maintain the experts profile [2]. The expertise knowledge management usually takes lots of effort and it becomes a hot topic to improve this task by using the outer source, such as web information [6, 8, 9]. However, we only have very short time to assign the reviewer for a proposal in our scenario. We focus on the problem of expert finding and expertise knowledge management in proposed reviewer recommendation system.

Our approach divides the problem of expert-finding into three parts. First, it reduces this problem into a multi-domain classification issue since we want to improve complexity of finding the experts and efficiency of the recommendation system. Second, it uses the *WCN* as a knowledge inference database and computes the correlated relatedness between experts and proposal. Finally, it takes account of the contribution in academic of each expert who belongs to the domain of proposal. It considers that includes user experience, research related, and authority of academic for reviewer recommendation system in real world.

1.1 Problem Definition

In this paper, we are required by the Division of Computer Science of National Science Council in Taiwan (NSC) to help the reviewer recommendation committee and find out suitable reviewers for research proposals. The dataset of NSC contains 8 domains, 668 proposals, 2,701 experts¹ who have one/more expertise domains, 38468 publications, 71,899 keywords of publication, and lots of submitted proposals for every year (i.e., 668 in this year for testing). For definition, each proposal is denoted as Pro_i and i is the index of proposals. Each expert which has a set of publications to represent the concept of his expertise is denoted as $Expert_k$, k is the index of expert. Each publication is denoted as Pub , and each $Expert_k$ publishes the Pub_{kj} , where j is the index of publication. Furthermore, Each $Expert_k$ has one/more domains and each Pro_i only belongs to a domain.

2 System Architecture

The system architecture of our approach comprises three parts: 1) Domain Modeling, 2) Expert Matching, and 3) Ranking, as shown in Figure 1. In this system, Domain Modeling improves the cost of computation and handles the problem of Expertise Knowledge Management. In the phase of Expert Finding, it solves the problem of correlation ranking by Expert Matching and estimates academic contributions of experts by Ranking.

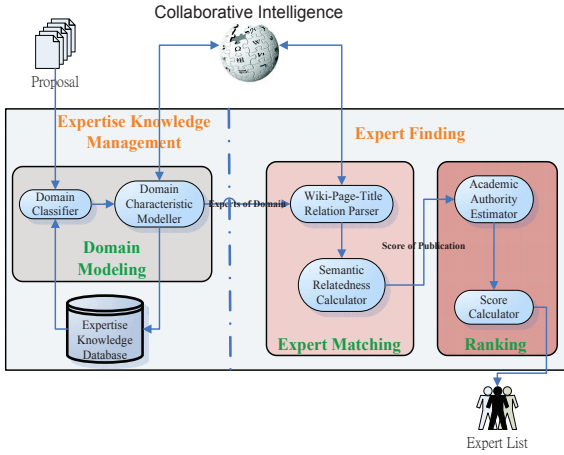


Figure 1. Reviewer recommendation system architecture.

2.1 Domain Modeling

The goal of Domain Modeling is to find relevant experts quick and generate domain knowledge efficiently. Building the *EKD* is necessary and helpful for find the suitable domain efficiently. Domain Characteristic Modeller is an incremental *EKD* learner for modeling the specific domain. Each expert has expertise domains and a set of Wikipedia Page Title (*WPT*) mapping from Wikipedia [9] to represent his research topic. Traditionally, finding the relevant experts usually costs a lot of time to compute the correlation with irrelevant experts. Hence, the proposed system classify the queried proposal first instead of comparing with all candidates for expert. The domain which we want to identify is notated as $DPro_n$, where n is the index of the domain. The set of *WPT* represents the concept of proposal and each *WPT* is denoted as $ProP_{iu}$, where i is the index of proposal and u is the index of *WPT*. This modulo classifies the proposals according to the probabilities of *WPT* for domains. After the classification, the unseen terms would be labeled as the concept. A function sums up the probability of *WPT* in the domains and calculates the probabilities for each domain. Another function calculates the probabilities by invoking Bayes' Theorem. It models the knowledge of a domain from associated documents. The probability function is as follows:

$$p(DPro_n|Pro_i) = \frac{p(Pro_i|DPro_n)p(DPro_n)}{p(Pro_i)} \quad (1)$$

For the purpose of finding the related domain, we assume the probability $p(Pro_i)$ to be uniform. We focus on $p(Pro_i|DPro_n)$ and $p(DPro_n)$. We estimate the probability of a proposal given a domain by representing the domain as a multinomial probability distribution over the keywords of proposal:

$$p(Pro_i|DPro_n) = \prod_{ProP_{iu} \in Pro_i} p(ProP_{iu}|\Theta DPro_n)^{n(ProP_{iu}, Pro_i)} \quad (2)$$

Then, we smooth the probability of a proposal's keyword given a domain with the background probabilities:

$$p(ProP_{iu}|\Theta DPro_n) = (1 - \lambda) * p(ProP_{iu}|DPro_n) + \lambda * p(ProP_{iu}) \quad (3)$$

where $\lambda = (\tau)/(\alpha + \beta)$, α is the average number of keywords in publication, and β is the average length of the publication title. And, $\alpha = 1.864$, $\beta = 60.112$, and the probability of $p(DPro_n)$ is computed by the count of proposals in domain n divides by the count of all the queried proposals. The both functions identify the domain which has the maximum probability as the answer of Pro_i .

¹Experts' data is retrieved from <http://cs.nsc.ncku.edu.tw/introduce/>

2.2 Expert Matching

The goal of Expert Matching is to measure the semantic relatedness between proposal of expert and publications. Wiki-Page-Title Relation Parser parses the Wikipedia categories of page as a concept. Since we want to measure the concept of relation between publication and proposal, the relation of terms should be measured first. The distances between each pair of categories are the degree of relation in *WCN*, and each pair have a maximum depth from root. According to the previous research, we take account of distance which limits to 5 because there is no relation and this pair of keyword would not be consider [6].

The relations of keywords are found, hence the score of semantic relatedness can be measured. The estimating criterion is that using the concept structure based on collaborative intelligence and finding the semantic relatedness in the Wikipedia. The *WCN* is a collaborative tagging system allowing users to categorize the content of page. The meaning of categories are from top to down, so the more specific concept it is, the more deeper depth of categories are. The distance between categories means the correlation, so the more correlated it is, the more closer distance of categories are. There are many pairs of keyword between proposal and publications and every pair has many paths to connect each other. The scores of each pair are computed by considering distance and depth, and the maximum one of these scores represents degree of semantic relatedness of this pair. Finally, sum of the maximum score of each pair is the score for measuring the semantic relatedness between proposal and publication.

2.3 Ranking

The goal of Ranking is to combine the scores of publications for each expert and rank the experts in the output list. The academic contribution can be estimated by number of publications. It computes $FinalScore_{Pro_i, Expert_k}$ of $Expert_k$ for Pro_i and the function is as follows:

$$FinalScore_{Pro_i, Expert_k} = \sum_{Pub_{kj} \in Expert_k} Score_{Pro_i, Pub_{kj}} \quad (4)$$

3 Experiments

In this paper, we examine two experiments on Domain Modeling and Expert Matching. The all combinations of features and methodologies are examined to find the best result for each domain. The domains include the domains of “Image and Pattern Recognition” (*IPR*), “Natural Language and Speech Processing” (*NLSP*), “Artificial Intelligence” (*AI*), “Computer Graphics” (*CG*), “Information

System Management” (*ISM*), “Database” (*DB*), “Bioinformatics” (*Bio*), and “Web Technologies” (*WT*).

3.1 Performance Analysis of Classifier

The performance of domain classifier depended on the methodology of classification and how to model the characteristics of domains, but there were lots of expert’s data to model the characteristic of domains. Hence, we wanted to use the fewer instance modeling domains, the features could be selected to examine, such as the keywords that were key-terms of proposal (*KT*), the title of Wikipedia page (*PT*), and the title of Wikipedia category (*CT*). The factors of multiple domains for modeling were adding weighting (*W*) or not (*NW*), and two methods of classification are max probability (*Max_Prob*) and Naïve Bayesian (*Naïve*). In

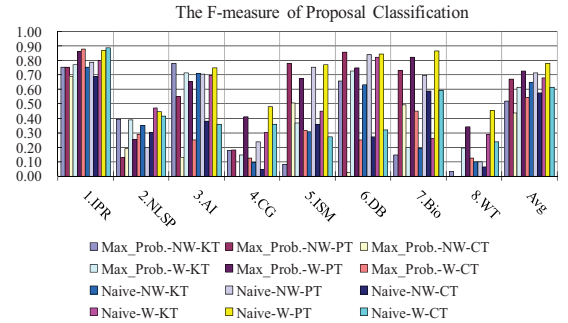


Figure 2. The F-measure of proposal classification.

this experiment, the performance of domain classifier is represented by the f-measure illustrated in Figure 2 against each domain. According to the result, *Naïve-W-PT* is better than other combinations of methodologies about 5% to 25% especially in the domain of *IPR*, *DB*, and *Bio*.

3.2 Correct Rate of Reviewer Recommendation

The criterion of evaluation is one of domains that expert had match the domain of proposal, and a labeled proposal corresponds to one domain. In our approach, experts of the domain for the proposal are ranked by Wikipedia relatedness score. The result shown in Figure 3 is compared with previous work. The expert data which is extracted from NSC website and these data are different from the data used by previous work. However, it is up-to-date and more fair than previous one. We wanted to make the top of expert list correct, but the answer of ranked expert list was not existent. Therefore, we evaluate the result by precision rate at

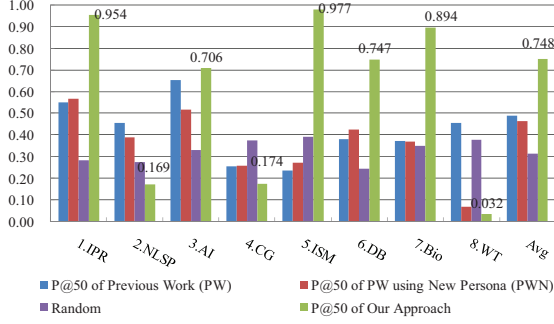


Figure 3. Precision@50 of related expert list.

$N(P@N)$, where N is the number of top result. The function is as following:

$$Precision(N, d) = \frac{\sum_{i \in Domain_d} \sum_{j=1}^N F(Result_{ij})}{PC_d} \quad (5)$$

where $F(Result_{ij})$ is $\frac{1}{j}$, if expert's domains match proposal's, else it is zero. d is the index of domain, PC_d is the proposal's count of domain d , i is the index of proposal, and j is the index of expert in the result list. The $P@50$ of PWN is almost equal to the PW , except for the domain of WT , hence we can compare with the result of PWN instead of the result of PW . The precision of previous work is equal to the result of random test in many domains, such as the domains of CG , ISM , Bio , and WT . The average $P@50$ of our approach is better than random about 40% and is better than PWN about 20%.

4 Conclusion

In this paper, we propose a reviewer recommendation system which assists the commissioner of national science organization to find the experts who are suitable to review the proposal. Our proposed system uses the publications of experts as the training data for the expertise knowledge database. The keywords of publications are transformed into domain concept from collaborative intelligence and the correlations between experts and proposal are considered by parsing WCN . The correlation which is between experts and proposal and the efficiency of the computation times are improved in our proposed approach, and the performances of convinced evaluation are shown in our experiments. The F-Measure of domain classifier is about 78.2% and $P@50$ of recommendation expert list is better than our previous work at least 20%. According to the result, many proposals have been classified into nothing domain which means we need more effort to complete the labeled terms in EKD . In

addition, the academic contribution are not considered very well, and it causes that the quantity of publications is more important than quality of publications.

Acknowledgements

This work was supported in part by the National Science Council of Taiwan under grants NSC 95-2221-E-001-021-MY3 and NSC 96-2628-E-011-084-MY3.

References

- [1] K. Balog, M. de Rijke, and W. Weerkamp. Bloggers as experts: feed distillation using expert retrieval models. In *SIGIR '08: Proceedings of the 31st annual international ACM SIGIR conference on Research and development in information retrieval*, pages 753–754, New York, NY, USA, 2008. ACM.
- [2] G. Demartini. Finding experts using wikipedia. *FEWS 2007: Finding Experts on the Web with Semantics Workshop at ISWC 2007 + ASWC 2007*, November 2007.
- [3] P. Liu, K. Liu, and J. Liu. Ontology-based expertise matching system within academia. pages 5431–5434, Sept. 2007.
- [4] M. Mika, M. Ciaramita, H. Zaragoza, and J. Atserias. Learning to tag and tagging to learn: A case study on wikipedia. *Intelligent Systems, IEEE*, 23(5):26–33, Sept.-Oct. 2008.
- [5] D. Mimno and A. McCallum. Expertise modeling for matching papers with reviewers. In *KDD '07: Proceedings of the 13th ACM SIGKDD international conference on Knowledge discovery and data mining*, pages 500–509, New York, NY, USA, 2007. ACM.
- [6] S. P. Ponzetto and M. Strube. Knowledge derived from wikipedia for computing semantic relatedness. *Journal of Artificial Intelligence Research*, 30:181–212, 2007.
- [7] F. Wang, B. Chen, and Z. Miao. A survey on reviewer assignment problem. In *IEA/AIE '08: Proceedings of the 21st international conference on Industrial, Engineering and Other Applications of Applied Intelligent Systems*, pages 718–727, Berlin, Heidelberg, 2008. Springer-Verlag.
- [8] P. Wang and C. Domeniconi. Building semantic kernels for text classification using wikipedia. In *KDD '08: Proceeding of the 14th ACM SIGKDD international conference on Knowledge discovery and data mining*, pages 713–721, New York, NY, USA, 2008. ACM.
- [9] K.-H. Yang, C.-Y. Chen, H.-M. Lee, and J.-M. Ho. Efs:expert finding system based on wikipedia link pattern analysis. In *the 2008 IEEE International Conference on Systems, Man and Cybernetics (SMC 2008)*, October 2008.
- [10] D. Yimam. Expert finding systems for organizations: Domain analysis and the demoir approach. In *ECSCW 99 Beyond Knowledge Management: Management Expertise Workshop*, pages 276–283. MIT Press, 2000.
- [11] J. Zhang, M. S. Ackerman, and L. Adamic. Expertise networks in online communities: structure and algorithms. In *WWW '07: Proceedings of the 16th international conference on World Wide Web*, pages 221–230, New York, NY, USA, 2007. ACM.