

Context-Aware Users' Preference Models by Integrating Real and Supposed Situation Data

Chihiro ONO^{†a)}, Yasuhiro TAKISHIMA[†], Yoichi MOTOMURA^{††}, Hideki ASOH^{†††}, *Members*,
Yasuhide SHINAGAWA^{††††}, *Nonmember*, Michita IMAI^{†††††}, and Yuichiro ANZAI^{†††††}, *Members*

SUMMARY This paper proposes a novel approach of constructing statistical preference models for context-aware personalized applications such as recommender systems. In constructing context-aware statistical preference models, one of the most important but difficult problems is acquiring a large amount of training data in various contexts/situations. In particular, some situations require a heavy workload to set them up or to collect subjects capable of answering the inquiries under those situations. Because of this difficulty, it is usually done to simply collect a small amount of data in a real situation, or to collect a large amount of data in a supposed situation, i.e., a situation that the subject pretends that he is in the specific situation to answer inquiries. However, both approaches have problems. As for the former approach, the performance of the constructed preference model is likely to be poor because the amount of data is small. For the latter approach, the data acquired in the supposed situation may differ from that acquired in the real situation. Nevertheless, the difference has not been taken seriously in existing researches. In this paper we propose methods of obtaining a better preference model by integrating a small amount of real situation data with a large amount of supposed situation data. The methods are evaluated using data regarding food preferences. The experimental results show that the precision of the preference model can be improved significantly.

key words: user preference model, context-awareness, recommender systems, probabilistic modeling

1. Introduction

Modeling users' preferences is a key technology for various personalized applications, such as recommender systems [1], [12], intelligent user interface, and one-to-one marketing. In recent years, "context-awareness" has become one of the most important research issues when constructing preference models. One of the reasons could be the fact that the diversification of contexts/situations in which the user uses the service, e.g. in town or in the home, as well as the diversification of the services and related items, has rockets together with the surge in Internet access via PDAs and cellular phones. These trends revealed that users' preferences drastically change based on the context/situation. For exam-

ple, movie preferences may change depending on the mood or the accompanied persons of the users.

Ono et al. have proposed a novel way to construct context-aware preference models using Bayesian networks [8], [9] and have implemented a movie recommender system on cellular phones using the model [10]. There, we have represented complex relations among users' profiles, contents' attributes, and situation attributes with a Bayesian network.

In constructing such context-aware statistical preference models, one of the most important but difficult problems is acquiring considerable training data in various situations. In particular, some situations require a heavy workload to set them up or to collect subjects capable of answering the inquiries involved. As an example, let consider a food recommender system that recommend food such as curry rice and beef bowl through cellular phone display by using users' information. The system should treat such conditions as "a user is choosing food in hot weather when feeling tired and hungry". In order to collect training data for such situation, the model constructor should make subjects tired and hungry, and gather them on a hot day.

As a way of reducing this difficulty, it is usually done to simply collect a small amount of data in a real situation, or to collect data in a supposed situation, i.e., a situation that the subject pretends that he is in the specific situation to answer inquiries, instead of setting up a real situation and putting them into it. Collecting answers in the supposed situations require much lighter workload than setting up a real situation. Although the data acquired in the supposed situation may differ from that acquired in a real situation, the difference is not taken seriously in existing researches and usually neglected. However, as we show in this paper, the difference is not negligible in some cases and the way of compensating the difference is necessary.

This paper is the first trial to compensate for the difference. The main idea is combining the data taken in supposed situations and that taken in real situations. In particular, we propose several methods of obtaining a better preference model by integrating small amount of real situation data and large amount of supposed situation data, and evaluate them using the data regarding food preferences.

The rest of this paper is organized as follows. Section 2 formulates the problem and related works are described in Sect. 3. Section 4 describes solutions, while experiment in food recommendation is shown in Sect. 5. Section 6 is for

Manuscript received April 8, 2008.

Manuscript revised July 4, 2008.

[†]The author is with KDDI R&D Labs, Inc., Fujimino-shi, 356-8502 Japan.

^{††}The author is with AIST, Tokyo, 135-0064 Japan.

^{†††}The author is with AIST, Tsukuba-shi, 305-8568 Japan

^{††††}The author is with West Japan Railway Company, Osaka-shi, 530-8341 Japan.

^{†††††}The author is with Keio University, Yokohama-shi, 223-0061 Japan.

a) E-mail: ono@kddilabs.jp

DOI: 10.1093/ietisy/e91-d.11.2552

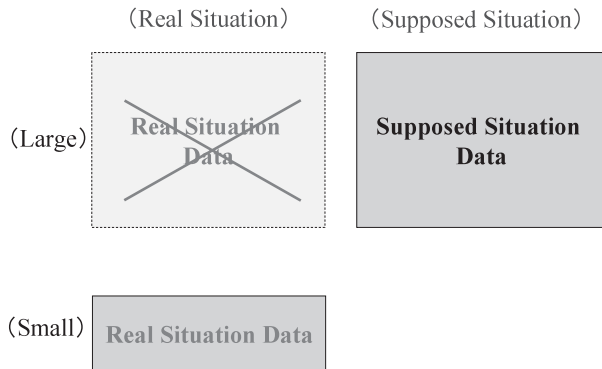


Fig. 1 Data acquisition.

evaluation of the models, while Sect. 7 describes conclusion and future works.

2. Problem

In this section we will formalize the problem treated in this paper.

2.1 Statistical Preference Model

There has been an abundance of research on constructing statistical preference models [15]. In most of them, the problem of constructing the statistical preference model is formalized as modeling the conditional probability distribution $P(V | U, C, S)$ to predict the value V from U , C , and S . Here, U represents a set of users' attribute variables, such as age, sex, etc. C represents a set of target contents' attribute variables, for example, category, calorie, fat, salt, etc. for foods. S represents a set of user situation/context variables such as hungry, tired etc. for food recommendations, and V denotes the user's preference/rating of a given content within a given context.

2.2 Issue

In order to construct context-aware statistical preference models, we should acquire considerable amount of training data in various situations. However, as is described above, it is often difficult to collect data under real conditions, because we have to set up a real situation and collect subjects capable of answering the question in the real situations.

In order to reduce the difficulty, we take a novel approach of collecting and using a small amount of real situation data and a large amount of supposed situation data instead of a large amount of real situation data as depicted in Fig. 1. Here the problem is how to get a better performance model by combining the real situation data and supposed situation data. We propose several integration methods and evaluate them using the data regarding food preferences.

3. Related Works

In the context of preference modeling and recommender

systems [1], [12], we have not found any research treating the same problem as ours to date. In most cases, this problem is not recognized and preference models are constructed with large amount of data taken by surveys in the supposed situation. Or very small amount of real situation data are used to construct preference models and result in the insufficient accuracy of the models.

In the broader context of statistical modeling, some related works do exist, of which the most relevant concern language model construction in the field of speech recognition, text mining, and the information retrieval. In those fields, statistical models of natural language, such as the n-gram model, are constructed using large scale natural language corpus [4], and the adaptation or modification of the statistical model has become a hot research issue recently. For example, in the text mining area, a language model adapted for texts describing a specific topic is created by modifying the general model, which is constructed with a large amount of data, with a small amount of text data describing the topic. In particular, some model adaptation methods such as interpolating statistical models are proposed in researches on the language model adaptation [6], [14].

The methods of updating statistics in non stationary conditions such as RLS (recursive least square) filter is also related to our problem. RLS has been applied to various problems including image tracking [15].

The contribution of our work is to try to generalize the ideas in those researches, apply them to the problem of compensating the real and supposed difference in preference model construction, and demonstrate the effectiveness.

4. Solution

4.1 Approaches

In order to solve the above problem, three approaches can be considered:

1. Data modification approach: To modify supposed situation data using real situation data. As an example, a large amount of supposed situation data are merged with or weighted by a small amount of real situation data. Subsequently, a model is constructed using the merged/weighted data.
2. Model modification approach: To modify and/or integrate a supposed situation model parameters with a real situation model parameters. Firstly, a base model is constructed using a large amount of supposed situation data. Subsequently, the model is modified in order to adapt to real situations with a small amount of real situation data.
3. Inference result modification approach: To modify and/or integrate an inference result of a supposed situation model with an inference result of a real situation model. Firstly, two models are constructed using a large amount of supposed situation data, and a small amount of real situation data, respectively. Sub-

sequently, the predicted preference value from the supposed situation model is modified and/or integrated with the value from the real situation model.

In the rest of this paper, we mainly focus on the second and third approaches. The reason is that, reconstructing the preference model with large amount of merged/weighted data every time is time consuming as it will be shown in 6.2. When operating on-line application systems of the preference model such as recommender systems, additional data in a real situation can be frequently acquired through, for example, user feedback or usage history. Here, fast model update against additional data is a very important issue. By the first approach, merging/re-weighting a large amount of data and reconstructing an updated model using them are necessary. On the contrary, in the second and third approaches, simply reconstructing the adapted model using a small amount of data is enough. The former procedure costs much more than the latter, and is undesirable for on-line applications.

4.2 Methods

In the following we will propose two methods corresponding to the second and the third approach. Both involve models being constructed using a large amount of supposed situation data and a small amount of real situation data. The inputs of the models are user attributes, situation attributes, and the attributes of the target content, and an output of the model represents user's preference of the content.

First, a model is constructed with the supposed situation data. For usual statistical models such as liner regression model, neural networks, and Bayesian networks, the model construction procedure includes two parts, model structure selection and model parameter estimation. We call the model constructed with the supposed situation data as a base model or supposed situation model, and denote the output value as V_s .

Subsequently, using the same model structure as the base model, model parameters are re-estimated with a small amount of real situation data. We call this model the real situation model and denote the output value as V_r .

4.2.1 Method 1: Model Modification

Due to the small amount of real situation data, the re-estimation process of the real situation model tends to become unstable. To relieve this difficulty, we propose smoothing the re-estimated parameters with parameters in the supposed situation model (See Fig. 3).

For example, if the model is represented as a Bayesian network, each CPT (Conditional Probability Table) in the network is smoothed by taking the weighted sum of the CPT in the real situation model and the CPT in the supposed situation model as:

$$P(x|pa(x)) = \alpha Pr(x|pa(x)) + (1 - \alpha)Ps(x|pa(x))$$

Here, $Pr(x|pa(x))$ denotes the CPT attached to variable x

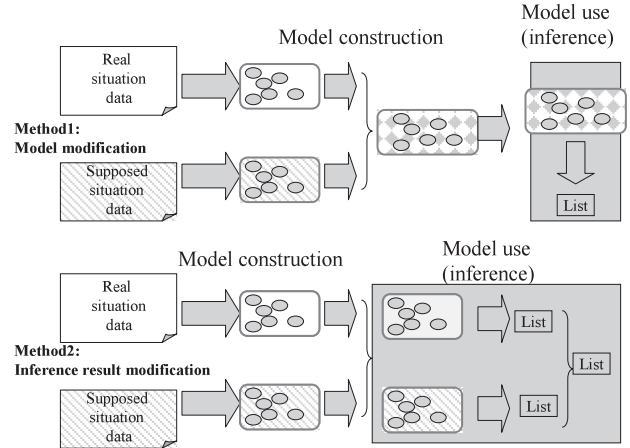


Fig. 2 Method 1 and method 2.

in the real situation model. $pa(x)$ is the set of parent nodes of x . $Ps(x|pa(x))$ denotes the corresponding CPT in the supposed model. α is a weight value.

This procedure of taking weighted sum of statistics itself is a very general technique and is used in some areas such as language model adaptation [6], [14] as we described in Sect. 3. It can be considered as a kind of technique for updating statistics such as RLS. RLS is also applied to various areas such as visual tracking [15].

4.2.2 Method 2: Inference Result Modification

This procedure is similar to the previous one. Here, instead of smoothing parameters in the model, we propose simply taking the weighted sum of two models (See Fig. 2). That is, the weighted sum of the output from the real situation model (V_r) and the output from the virtual situation model (V_s) becomes the predicted evaluation:

$$\hat{V} = \alpha V_r + (1 - \alpha) V_s$$

When the model is linear relative to the parameters, it becomes equivalent to the method 1. However, if the model is non-linear, such as in a Bayesian network, then the result differs.

There are several approaches to improve the performance of statistical models by combining multiple classifiers. For example, “Bagging” is a kind of meta-classifier which divides training set, and generates many models based on divided data sets, and estimates final decision by these voting [2]. In terms of combining results from multiple classifiers, this method looks similar to existing meta-classifier techniques. However, the existing meta-classifier techniques cannot be directly applied to our problem because the condition of our problem is different from those of meta-classifier techniques. That is, existing meta-classifier techniques including bagging applied the idea of combining multiple classifiers to a training data gathered in a “single situation” in order to get more robust/accurate classifier for the situation. To the contrary, here we have applied the idea of combining multiple classifiers to different training data

gathered in a “two different situations”, the supposed situation data and the real situation data, in order to get the better classifier for the “real situation”. Thus, the voting technique used in bagging scheme, which votes for the multiple outputs from the same kind of data, cannot be applied.

5. Experiment in Food Recommendation

We evaluated the effectiveness of the proposed methods in a food recommender system which we are developing.

5.1 Food Recommender System

We are developing a food recommender system for cellular phone users. In the system, a Bayesian network is used to model the joint probability distribution $P(V, U, C, S)$. As is well known, in a Bayesian network, each random variable is represented as a network node, and the network links represent dependencies between variables. Conditional independences between variables are represented by the entire network structure and used for a more efficient probabilistic inference [5], [7], [11].

Figure 3 shows the overview of the recommender system, while Fig. 4 shows the flow of a recommendation process. Firstly, a user sends a request for recommendation based on the situation attributes (degree of hunger, degree of fatigue, and daily temperature) through his/her cellular phone. Subsequently, the recommender system merges the registered user attributes with the input user situational attributes, calculates the probability of the user rating for each candidate food using the Bayesian network inference engine, and composes a recommendation list of foods according to the probability of positive ratings (See Fig. 5). The recommendation system may receive user feedback, and periodically, the system updates the parameters of the Bayesian network model using feedback data by using the Bayesian inference engine in order to increase the precision of the recommendation.

5.2 Data Acquisition for Model Construction

The data acquisition procedure was composed of two parts: a large-scale virtual-situation questionnaire survey and a small-scale real/virtual situation questionnaire survey.

Firstly, a questionnaire survey concerning supposed situations was conducted in December 2006.

- Number of subjects: 1500
- Number of foods: 20
- Inquiries:
- User demographic and lifestyle attributes: 47 attributes such as age, gender, and occupation, brand loyalty, time and expenditure on leisure.
- User attributes regarding food appreciation: 19 attributes such as food category preference. (3-grade scale for each attribute)
- Rating of the food and reasons of the rating for each

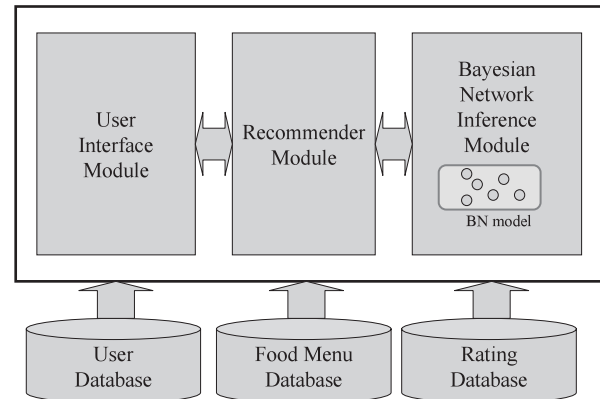


Fig. 3 Overview of recommender system.

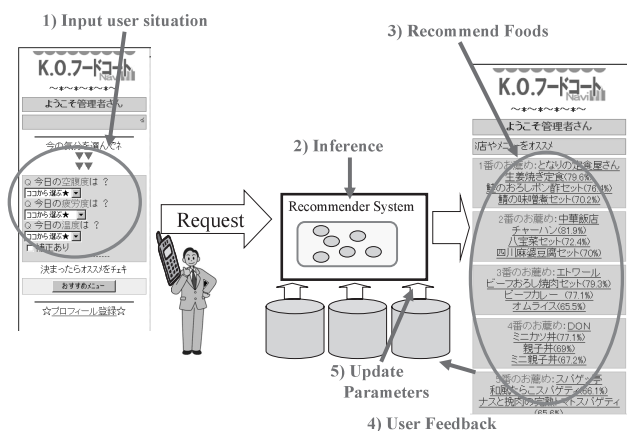


Fig. 4 Flow of recommendation process.



Fig. 5 Screenshot of cellular phone display.

food under the given situations in Table 1: 1 total evaluation (3-grade scale from satisfied very much to not at all satisfied) and 19 reasons (impressions of the food such as salty, easy to eat, etc.) (3-grade scale for each reason)

Each subject rated food 14–16 times: in 3–4 situations and 3–4 foods for each situation.

Table 1 Situations.

ID	Situation
S0	Neutral condition
S1	Starving condition: e.g. before dinner without lunch.
S2	Almost full condition: e.g. 1 to 2 hours after dinner.
S3	Very tired condition: e.g. right after tough sports.
S4	Slightly tired condition: e.g. after jogging.
S5	Hot weather condition
S6	Cold weather condition

Table 2 Data sets.

	Number of Training Data	Number of Test Data
Supposed Situation Data	24000 (=23760+240)	0
Real Situation Data	240	240

Consequently, we obtained 23760 records of supposed situation data, we call them dataset 1.

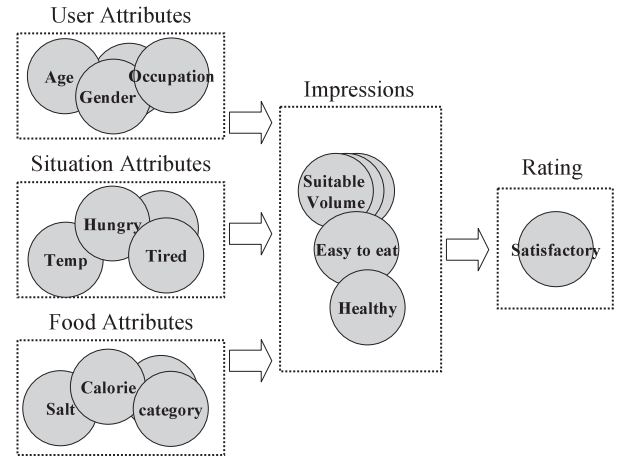
Secondly, a questionnaire survey concerning both real and supposed situations was conducted in March 2007.

- Number of subjects: 12
- Number of foods: 20
- Inquiries: the same as the first questionnaire. Each subject rated 4 to 20 foods under situations in Table 1.

Here, we obtained 480 records of real situation data and 240 records of supposed situation data. We call them dataset 2. We divided the real situation data in dataset 2 into two parts evenly: training data for model construction and test data for model evaluation. The supposed situation data in dataset 2 are merged with the dataset 1. A total of 24000 records were used for constructing the supposed situation model. The 240 training data from the real situation data, which is set to 1/100 of the supposed situation data, were used to construct the real situation model. The 240 test data from real situation data were reserved for evaluation of the integration methods. Detail is shown in Table 2.

5.3 Construction of Base Model

Due to the considerable number of random variables in our model, we specify the rough network structure manually and estimate the conditional probabilities from the data. We initially assume a rough network structure used to predict the user's ratings for foods, which is shown in Fig. 6. This structure means that the overall rating depends on common variables representing the user's impressions of a food (suitable seasoning, suitable volume, etc.), and the impressions depends on user attributes, situational attributes, and food attributes. As food attributes, we use 8 attributes such as calories, salt.

**Fig. 6** Model structure.

For the model used in the following experiments, we reverse the direction of links in order to simplify the CPT. We selected effective variables (with the maximum number of parent nodes set to 5) based on many observed attributes, determined local network structures reflecting the relationship between selected variables, and estimated the CPTs via a standard maximum likelihood estimation, using the 24000 supposed situation data. (For more detailed model construction steps, please see [8]).

5.4 Model Modification by Method 1

The above base (supposed situation) model is modified according to methods 1 in Sect. 3. Here, the CPT of each node is re-estimated using the 240 real situation data. Subsequently, each CPT from the base model and real situation model are merged by taking the weighted sum of the corresponding probability values.

5.5 Inference Result Modification by Method 2

In method 2, after constructing real situation model by re-estimating CPT using the 240 real situation data, the output of base (supposed situation) model and real situation model are merged by taking the weighted sum of the estimated evaluation values.

6. Evaluation

The constructed models are evaluated according to the accuracy of their preference predictions.

6.1 Evaluation Criteria

As a measure of accuracy, we used the mean squared error (MSE) of the prediction. When the total number of predicted ratings is N , the number of values of the rating is r , the correct rating value of User i to Food j in Situation k is p_{ijk} , the predicted rating value is v , the MSE can be formulated as:

Table 3 Results of method1.

α	Mean	Std Dev.
0	0.724	0.038
0.1	0.664	0.039
0.2	0.647	0.051
0.3	0.662	0.065
0.4	0.687	0.074
0.5	0.714	0.081
0.6	0.74	0.089
0.7	0.766	0.096
0.8	0.794	0.105
0.9	0.825	0.112
1	0.855	0.117

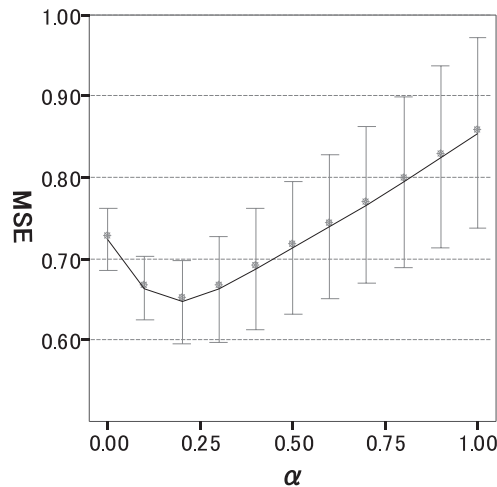
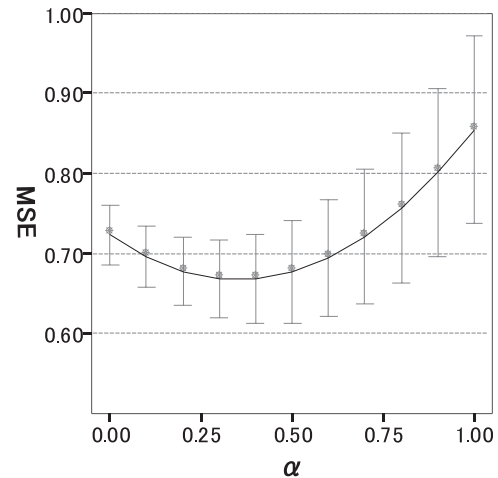

Fig. 7 MSE values for method 1.

Table 4 Results of method 2.

α	Mean	Std Dev.
0	0.723	0.037
0.1	0.696	0.038
0.2	0.678	0.042
0.3	0.668	0.048
0.4	0.668	0.056
0.5	0.676	0.064
0.6	0.694	0.073
0.7	0.721	0.083
0.8	0.757	0.094
0.9	0.801	0.105
1	0.855	0.117


Fig. 8 MSE values for method 2.

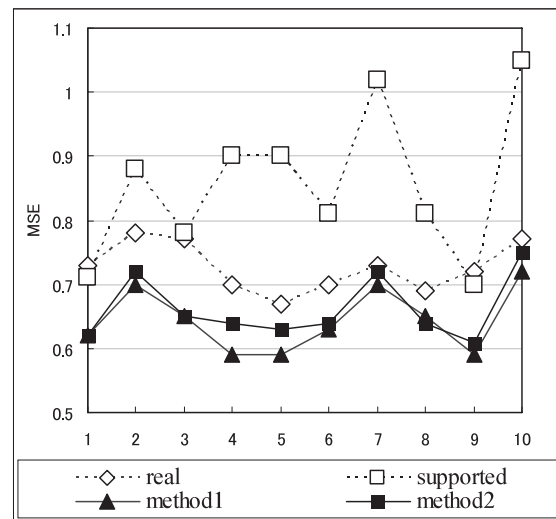
$$\frac{1}{N} \sum_{ijk} \left(p_{ijk} - \sum_{v=1}^r v P(V=v | U=i, C=j, S=k) \right)^2.$$

6.2 Results

We repeated the experiments 10 times using different divisions of training data and test data sets. Table 3 and Fig. 7 show the results for method 1, and Table 4 and Fig. 8 show the results for method 2, which includes the mean and standard deviation of MSE scores against 240 real-situation test data when the weight ($= \alpha$ in 4.2.1 and 4.2.2) moves from 0 to 1. Figure 9 shows the best MSE scores attained by the best weighting value α for each of 10 trials. We have applied the paired t test to compare the MSE values of $\alpha = 0$ and $\alpha = 0.2$ for the method 1, $\alpha = 0$ and $\alpha = 0.4$ for the method 2. In both cases, it was confirmed that the improvements of the MSE scores are significant even at the 1% confidence level.

The evaluation results illustrate that the performance of the base model constructed from supposed situation data is superior to the real model constructed from real situation data. This is because the volume of supposed situation data is much larger than that of real situations.

As illustrated in Table 3, 4 and Fig. 7, 8, 9, both


Fig. 9 Best MSE values for each trial.

methods 1 and 2 work well. In other words, the appropriately weighted model outperforms both the base model and the real situation model, which means that we can obtain preference model with higher accuracy than conventional approaches by integrating small amount of supposed

situation data and large amount of real situation data.

As a reference, we also evaluated the performance of a simple method of the data modification approach (approach 1 in Sect.4.1). We construct preference model with merged real and supposed situation data (24,240 records) and evaluated the accuracy and the computational time.

The mean and standard deviation of MSE scores against 240 real-situation test data is 0.689 and 0.040, respectively. This outperforms both the supposed situation model and the real situation model but worse than the best performance of method 1 (0.647) and method 2 (0.668).

In terms of the computational time of the model update, approach 1, approach 2 (method 1) and approach 3 (method 2) takes 1589 sec, 2 sec, 2 sec, at the same PC environment respectively. This result means that the approach 1 takes approximately 800 times as much time as that of the approach 2 and 3, and is not suitable for on-line applications as we wrote in 4.1.

7. Conclusion and Future Works

Motivated by the difficulty of data acquisition in real conditions, in this paper we formalized the problem of obtaining a better preference model by combining real situation data and supposed (imaged) situation data. We also proposed two methods to obtain a superior model and confirmed the performance improvement against existing approaches through the experiments regarding food preference.

As a future study, we would like to investigate the difference between the real situation and supposed situation in more depth, and propose more effective modeling methods. In particular, methods to estimate the optimal weight value α should be investigated.

The proposed methods can be applied to various kinds of statistical preference models other than Bayesian networks. In the area of statistical language modeling, various techniques such pLSA [3] and Dirichlet mixtures [13] have been applied to the model adaptation problem. Applying such advanced ideas to our problem may be an interesting issue.

From business viewpoints, getting better models with small costs is very important issue. If the amount of real situation data increases, the accuracy of models may become better. However the cost of survey becomes also high. There may be the optimal balance of real situation data and supposed situation data according to the budget of the survey. The method for finding the cost-effective combination is also a big issue.

Implementing applications such as recommender systems using the adapted model and evaluating the improvement of users' impressions to the services are also important future work.

By trying to apply to various targets other than food recommendation and refining our proposal to make it to be generic, we hope it opens up a new research area not only in preference modeling and intelligent systems using the

models, but also statistical modeling using survey data in general.

Acknowledgment

The authors wish to thank, Dr. Shigeyuki Akiba, President and CEO of KDDI R&D Laboratories, Inc. for his continuous support for this study.

References

- [1] G. Adomavicius and A. Tuzhilin, "Toward the next generation of recommender systems: A survey of the state-of-the-art and possible extensions," *IEEE Trans. Knowl. Data Eng.*, vol.17, no.6, pp.734–749, 2005.
- [2] L. Breiman, "Bagging predictors," *Mach. Learn.*, vol.24, no.2, pp.123–140, 1996.
- [3] D. Gildea and T. Hofmann, "Topic-based language models using EM," *Proc. 6th European Conference on Speech Communication and Technology (EUROSPEECH-99)* 5, pp.2167–2170, 1999.
- [4] X. Huang, A. Acero, and H.-W. Hon, "Language modeling," in *Spoken Language Processing: A Guide to Theory, Algorithm, and System Development*, Chapter 11, pp.545–590, Prentice Hall, 2001.
- [5] F.V. Jensen, *Bayesian Networks and Decision Graphs*, Springer-Verlag, 2001.
- [6] R. Kuhn and R. De. Mori, "A cache-based natural language model for speech recognition," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol.12, no.6, pp.570–583, 1990.
- [7] Y. Motomura, "Bayesian network construction system: BAYONET," *Proc. Tutorial on Bayesian Networks*, pp.54–58, 2001.
- [8] C. Ono, M. Kurokawa, Y. Motomura, and H. Asoh, "A. Context-aware movie preference model using a Bayesian network for recommendation and promotion," *Proc. User Modeling 2007 11th International Conference, UM 2007, Corfu, Greece, July, 2007, Proceedings, Lecture Notes in Computer Science 4511*, pp.247–257, 2007.
- [9] C. Ono, Y. Motomura, and H. Asoh, "A study of probabilistic models for integrating collaborative and content-based recommendation," *Working Notes of IJCAI-05 Workshop on Advances in Preference Handling*, 2005.
- [10] C. Ono, M. Kurokawa, Y. Motomura, and H. Asoh, "Implementation and evaluation of a movie recommender system considering both users' personality and situation," *IPSJ J.*, vol.49, no.1, pp.130–140, 2008.
- [11] J. Pearl, *Probabilistic Reasoning in Intelligent Systems: Networks of Plausible Inference*, Morgan Kaufmann Publishers, 1988.
- [12] P. Resnick and H.R. Varian, "Recommender systems," *Commun. ACM*, vol.40, no.3, pp.56–58, 1997.
- [13] K. Sadamitsu, T. Mishina, and M. Yamamoto, "Topic-based language models using Dirichlet mixtures," *System and Computers in Japan*, vol.38, pp.76–85, 2007.
- [14] K. Seymore and R. Rosenfeld, "Large-scale topic detection and language model adaptation," *Technical Report CMU-CS-97-152*, School of Computer Science, Carnegie Mellon University, 1997.
- [15] C. Wren, A. Azarbayejani, T. Darrell, and A. Pentland, "Pfnder: Real-time tracking of the human body," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol.19, no.7, pp.780–785, 1997.
- [16] I. Zekerman and D.W. Alberecht, "Predictive statistical models for user modeling," *User Modeling and User-Adapted Interaction*, vol.11, no.1-2, pp.5–18, 2001.



Chihiro Ono received the B.E. degree of Electrical Engineering, the M.S. degree of Computer Science from Keio University, Japan, in 1992 and 1994 respectively. Since joining KDD in 1994, he has been working on network management systems, database systems, and personalization technologies, from 1999 to 2000, he was a visiting researcher at Stanford University. He is currently a senior research engineer of Intelligent Media Processing Laboratory in KDDI R&D Laboratories, Inc. He received Best Paper Award for Young Researchers of the National Convention of IPSJ in 1996.



Yasuhiro Takishima received the B.S., M.S. and Ph. D in electrical engineering from the University of Tokyo, in 1986, 1988 and 1997 respectively. Since 1988 he has been with the Research and Development Laboratories of KDD and KDDI R&D Laboratories Inc., Japan, and engaged in research on the information theory and communication systems for such as versatile binary code design, high quality video compression, and image processing. His recent interests are intelligent media analysis and expression. Currently he is an senior manager in Intelligent Media Processing Laboratory.



Yoichi Motomura is a senior research scientist of Digital Human Research Center (DHRC) at National Institute of Advanced Industrial Science and Technology (AIST) in Japan. He is also the chief research officer of Modellize inc. He received M.S. from the University of electro-communications and then joined to the Real World Computing project in the electro technical laboratory, AIST, MITI in 1993. His current research works are statistical learning theory, probabilistic inference algorithms on Bayesian

networks, and their applications to user modeling and human behavior understanding. He received a best presentation award and research promotive award from Japanese Society of Artificial Intelligence.



Hideki Asoh is Senior Research Scientist in the Information Technology Research Institute, National Institute of Advanced Industrial Science and Technology. He received his Master of Engineering degree in Information Technology from the graduate school of the University of Tokyo in 1983. His research interests are in theory and application of statistical learning. He is a member of Japanese Society of Artificial Intelligence, the Japanese Society of Neural Networks, and the Behaviormetric Society of Japan.



Yasuhide Shinagawa received the B.S. and M.S. degrees in Engineering from Keio University in 2005 and 2007, respectively. He has been working on human robot interaction and semantic sensor network technologies. He is currently working at West Japan Railway Company.



Michita Imai is Associate Professor of Faculty of science and technology at Keio university and a Researcher at ATR Intelligent Robot Laboratories. He received his Ph.D. degree in Computer Science from Keio Univ. in 2002. In 1994, he joined NTT Human Interface Laboratories. He joined the ATR Media Integration & Communications Research Laboratories in 1997. His research interests include autonomous robots, human-robot interaction, speech dialogue systems, humanoid, and spontaneous behaviors. He

is a member of the Information Processing Society of Japan, the Japanese Cognitive Science Society, the Japanese Society for Artificial Intelligence, Human Interface Society, IEEE, and ACM.



Yuichiro Anzai is President of Keio University. He received B.E. in Applied Chemistry, M.E. and Ph.D in Administration Engineering from Keio University in 1969, 1971 and 1974, respectively. During 1979–1985 he was assistant professor, Faculty of Engineering at Keio University. During 1985–1988 he was Associate Professor, Faculty of Letters at Hokkaido University. Since 1988, he has been Professor, Faculty of Science and Technology at Keio University. During 1993–2001 he was Dean, Faculty of Science and Technology, Keio University. His research includes the methodological analysis of human cognition, thought, learning and memory processes, and the design of interactive systems between humans, their environment, computers and robots.