

Mining Fuzzy Association Rules Using Mutual Information

S. Lotfi, M.H. Sadreddini

Abstract— Quantitative Association Rule (QAR) mining has been recognized as an influential research problem over the last decade due to the popularity of quantitative databases and the usefulness of association rules in real life. However, the combination of these quantitative attributes and their value intervals always rise to the generation of an explosively large number of itemsets, thereby severely degrading the mining efficiency. In this paper, we introduce a novel technique, called MFAMI, for mining quantitative association rules using fuzzy set theory. MFAMI employs linguistic terms to represent the revealed regularities and exceptions. This algorithm avoids the costly generation of a large number of candidate sets. Instead, using mutual information indicates the strong informative relationships among the attributes; potential frequent itemsets will be discovered. By utilizing those itemsets we devise an efficient algorithm that significantly generates rules. For effective mine rules, MFAMI employs adjusted difference analysis with this advantage that it does not require any user-supplied thresholds which are often hard to determine. Since the proposed algorithm greatly reduces the candidate subsequence generation efforts, the performance is improved significantly. Experiments show that the proposed algorithm is capable of discovering meaningful and useful fuzzy association rules in an effective manner, speeding up the mining process and obtaining most of the high confidence QARs.

Index Terms—Quantitative Association Rules, Mutual Information, fuzzy theory, frequent itemset.

I. INTRODUCTION

Data mining, the effective discovery of correlations among the underlying data in large databases, has been recognized as an important area for database research and has also attracted a lot of attention from the industry as it has many applications in marketing, financial, and retail sectors. One commonly used representation to describe these correlations is called association rules as introduced in [1]. In this model, the set $I = \{i_1, i_2, \dots, i_m\}$ is a collection of items or attributes. The database DB consists of a set of transactions, where each transaction is a subset of items in I. An association rule is an implication of the form $X \Rightarrow Y$ with $X, Y \subset I$ and $X \cap Y = \emptyset$. The meaning of the rule is that a transaction containing items in X will likely contain items in Y. To determine whether an association rule is interesting, two measures are used: support and confidence. An association rule, $X \Rightarrow Y$, has support $s\%$ in DB if $s\%$ of transactions in DB contains items in $X \cup Y$. The same

association rule is said to have confidence $c\%$ if among the transactions containing items in X, there are $c\%$ of them containing also items in Y. So, the problem is to find all association rules which satisfy predefined minimum support and minimum confidence constraints.

In this setting, attributes which represent the items are assumed to have only two values and thus are referred as Boolean attributes. If an item is contained in a transaction, the corresponding attribute value will be 1; otherwise the value will be 0. Many interesting and efficient algorithms have been proposed for mining association rules for these Boolean attributes, for examples, Apriori [1], DHP [2], and partition algorithms [3]. However, in a real database, attributes can be quantitative and the corresponding domains can have multiple values or a continuous range of values, for examples, age, and salary. A common approach to the QAR mining problem is to transform it into a problem of conventional BAR mining [1],[4]. Existing algorithms (e.g. [4]-[6]) involve discretizing the domains of quantitative attributes into intervals so as to discover quantitative association rules. For each distinct value of a quantitative or categorical attribute, the pair <attribute, value> is mapped to a Boolean attribute and then algorithms for mining BARs are applied.

In many cases, the number of intervals associated with an attribute is large hence when we join the attributes in the mining process, the number of itemsets (i.e., a set of <attribute, interval> pairs) can become prohibitively large. As a result, effective techniques to prune the large search space of QAR mining and avoid the costly generation of a large number of candidate sets are necessary in order to develop an efficient algorithm for the problem. Also these intervals may not be meaningful enough for human experts to easily obtain nontrivial knowledge [7]. On the other hand, we can use fuzzy association rules which provide a smooth boundary, where each attribute will have a fuzzy set.

To extract fuzzy association rules, in this paper, information-theoretic measure will be adopted. Using this new measure, potential frequent itemsets will be discovered. By utilizing these frequent itemsets; a set of 1-dimensional rules are generated. This set is then filtered by adjusted difference measure. Combining the rules within this set in the next step, results in the set of 2-dimensional candidate rules and so on. These new measures, will lead to an efficient and scalable algorithm, MFAMI, for mining quantitative association rules.

The remaining parts of this paper are organized as follows: Related research is reviewed in Sect. 2. Then we give some preliminaries on QAR mining in Sect. 3. The proposed data-mining algorithm is described in Sect. 4. Experiments to demonstrate the performance of the proposed algorithm are stated in Sect. 5. Conclusions are finally given in Sect. 6.

Manuscript transmitted December, 2008 for ICDMA 2009.

S. Lotfi is with Department of Computer Science and Engineering, Shiraz University, Iran (e-mail: lsomayeh@gmail.com).

M. H. Sadreddini is with Department of Computer Science and Engineering, Shiraz University, Iran (e-mail: sadredin@shirazu.ac.ir).

II. RELATED WORKS

Quantitative association rule mining problem has been introduced in [5] and some algorithms for quantitative values also have been proposed, where the algorithm finds association rules by partitioning the attribute domain, combining adjacent partitions and then transforming the problem into a binary state.

Mining QARs by a generic BAR mining algorithm, however, is infeasible in most cases for the following reasons. First, QAR mining suffers from the same problem of a combinatorial explosion of attribute sets as does BAR mining; that is, given a set of N distinct attributes, the number of its non-empty subsets is $(2^N - 1)$. However, as shown by [5], it is necessary to combine the consecutive intervals of a quantitative attribute to gain sufficient support and more meaningful intervals. This leads to another combinatorial explosion problem: if the domain of a quantitative attribute is partitioned into n intervals, the total number of intervals of the attribute grows to $O(n^2)$ after combining the consecutive intervals. When we join the attributes in the mining process, the number of itemsets (i.e., a set of <attribute, interval> pairs) can become prohibitively large if the number of intervals associated with an attribute is large.

The second one is caused by the sharp boundary between intervals. To dominant this problem, Mining fuzzy association rules for quantitative values has been considered by a number of researches [8]-[13], most of which have based their methods on the important APriori algorithm. Chan and Au introduced F-APACS for mining fuzzy association rules [14]. Instead of using intervals, F-APACS employs linguistic terms to represent the revealed regularities and exceptions [15]. Kuok's algorithm [10] expects the user or an expert to provide the required fuzzy sets of the quantitative attributes and their corresponding membership functions. Fu argues that experts may not give the right fuzzy sets and their corresponding membership functions. Hence, he proposed a method to find the fuzzy sets based on clustering techniques [6]. Each of these researchers treated all attributes (or all the linguistic terms) as uniform. However, in real-world applications, the users perhaps have more interest in the rules that contain fashionable items. Gyenesi [7] introduces the problem of mining weighted quantitative association rules based on fuzzy approach. He assigns weights to the fuzzy sets to reflect their importance to the user and proposes two different definitions of weighted support: with and without normalization similar to his previous method.

Ishibuchi et al. extended the genetic algorithm-based rule selection method in Ref. [16] to the case where various fuzzy partitions with different granularities are used for each input. This extension increases the number of candidate rules. Hence, they proposed a prescreening procedure which is based on two rule evaluation criteria of association rules, to decrease the number of candidate rules. Kaya et al. [17] proposed an automated clustering method based on multi-objective genetic algorithms. This method automatically clusters the values of a given quantitative attribute in order to obtain large number of large itemsets in low duration. They compared their proposed approach with CURE-based approach. In addition to the autonomous

specification of fuzzy sets, experimental results exhibit good performance over CURE-based approach in terms of runtime as well as the number of large itemsets and interesting association rules.

In [22] to indicate the strong informative relationship among the attributes, a mutual information graph was constructed. The cliques in the MI graph represent a majority of the frequent itemsets. By utilizing the cliques in the MI graph, frequent itemset was computed. Frequent itemsets and prefix tree structure have been used to generate QARs. In this article the technologies of fuzzy sets and association rules were combined and extended. And then a fuzzy data mining algorithm was proposed to discover fuzzy association rules of weighted quantitative data. The discovered rules which are expressed in natural language are more understandable to human.

But finding all frequent itemsets in large databases with this algorithm requires multiple database scans, using complicated data structures that imposes extra space, computation and time overhead.

III. PRELIMINARIES

In this section, we present the notions and basic concepts in the QAR mining problem.

A. Definitions

- 1) Let $I = \{x_1, x_2, \dots, x_m\}$ be a set of distinct attributes where I_j ($1 \leq j \leq m$) shows the j -th item and m is the number of unique items. Attributes can be either quantitative or categorical. In item x [l_x, u_x], if x is categorical $l_x = u_x$ and $l_x \leq u_x$ if x is quantitative.
- 2) Let $D = \{T_1, T_2, \dots, T_n\}$ denote a quantitative database, where $T_i \subset I$ is called a *transaction*, n is the number of transactions. A transaction T is a sequence $\{v_1, v_2, \dots, v_m\}$, where $v_j \in \text{dom}(x_j)$, for $1 \leq j \leq m$.
- 3) A transaction T supports an itemset X if $\forall x_i [l_i, u_i] \in X, l_i \leq v_i \leq u_i$, where $i \in \{1, \dots, m\}$. The frequency of X in D , denoted by $\text{freq}(X)$, is the number of transactions in D that support X . The support of X , denoted by $\text{supp}(X)$, is the probability that a transaction T in D supports X , and is defined as $\text{supp}(X) = \text{freq}(X) / |D|$. X is a *frequent itemset* if $\text{supp}(X) \geq \sigma$, where σ ($0 \leq \sigma \leq 1$) is a predefined minimum support threshold (*minSup*).
- 4) A quantitative association rule (QAR), R , is an implication of the form $X \Rightarrow Y$, where X and Y are itemsets, and $\text{attr}(X) \cap \text{attr}(Y) = \emptyset$. X and Y are called the antecedent and the consequent of R , respectively. We define the attribute set of R as $\text{attr}(R) = \text{attr}(X) \cup \text{attr}(Y)$. The support of R is defined as $\text{supp}(X \cup Y)$. The confidence of R is defined as $\text{conf}(R) = \text{supp}(X \cup Y) / \text{supp}(X)$, which is the conditional probability that a transaction T supports Y , given that T supports X .
- 5) We predefine suitable linguistic terms (fuzzy regions) and their corresponding membership functions to transform numeric or categorical data into fuzzy values. Therefore assign some fuzzy set to each attributes, $R = \{R_{j1}, R_{j2}, \dots, R_{jk}\}$, such that R_{jk} is k -th fuzzy region of I_j . Membership degree of each quantitative or categorical value of I_j for T_i in fuzzy R_{jk} is $F_j^{(i)} = \{f_{j1}^{(i)}, f_{j2}^{(i)}, \dots, f_{jk}^{(i)}\}$.
- 6) Support of each fuzzy region R_{jk} , was computed using

formula 1.

$$Sup(R_{jk}) = \frac{\sum_{i=1}^n f_{jk}^{(i)}}{n} \quad (1)$$

- 7) For each candidate item S ($s_1, s_2, \dots, s_t, \dots, s_{t+1}$), support was computed using formula 2.

$$Sup(S) = \frac{\sum_{i=1}^n f_s^{(i)}}{n} = \frac{\sum_{t=1}^n \min_{t=1}^{r+1} f_{s_t}^{(i)}}{n} \quad (2)$$

- 8) Confidence value of each association rule $S_1 \wedge \dots \wedge S_x \wedge S_y \wedge \dots \wedge S_q \rightarrow S_k$, was calculated using formula 3.

$$Conf(S_1 \wedge \dots \wedge S_x \wedge S_y \wedge \dots \wedge S_q \rightarrow S_k) = \frac{\sum_{i=1}^n \min_{k=1}^q \omega_{s_k} f_{s_k}^{(i)}}{\sum_{t=1}^n \min(\min_{k=1}^x \omega_{s_k} f_{s_k}^{(i)}, \min_{k=y}^q \omega_{s_k} f_{s_k}^{(i)})} \quad (3)$$

B. Entropy and mutual information

Let x and y be two random variables. Given $v_x \in dom(x)$ and $v_y \in dom(y)$, we denote the probability parameters as follows:

- $p(v_x)$: the probability of x taking the value v_x .
- $p(v_x, v_y)$: the joint probability of x taking the value v_x and y taking the value v_y .
- $p(v_y | v_x)$: the conditional probability of y taking the value v_y given that x takes the value v_x . It is defined as $p(v_y | v_x) = p(v_x, v_y) / p(v_x)$.

In the QAR mining context, we have

- $p(v_x) = supp(x[v_x, v_x])$
- $p(v_x, v_y) = supp(x[v_x, v_x], y[v_y, v_y])$ and
- $p(v_y | v_x) = conf(x[v_x, v_x] \Rightarrow y[v_y, v_y])$.

Entropy: Entropy is a central notion in information theory [20], which measures the uncertainty in a random variable. Entropy and mutual information are closely related. The *entropy* of a random variable x , denoted as $H(x)$, is defined as:

$$H(x) = - \sum_{v_x \in dom(x)} p(v_x) \log p(v_x) \quad (4)$$

Mutual information: Mutual information describes how much information one random variable tells about another one. The *mutual information* of two random variables x and y , denoted as $I(x, y)$, is defined as:

$$I = \sum_{v_x \in dom(x)} \sum_{v_y \in dom(y)} p(v_x, v_y) \log \frac{p(v_x, v_y)}{p(v_x)p(v_y)} \quad (5)$$

The information that y tells us about x is the reduction in uncertainty about x due to the knowledge of y , and similarly for the information that x tells about y . the greater values of $I(x, y)$, the more information x and y tell about each other.

Normalized Mutual Information: The normalized mutual information of two attributes x and y , denoted as $\tilde{I}(x, y)$ is defined as:

$$\tilde{I}(x, y) = \frac{I(x, y)}{I(x, x)} \quad (6)$$

such that $I(x, x) = H(x)$.

IV. RULE CONSTRUCTION

To find quantitative association rules, in [19] a mining algorithm was proposed based on the concept of large itemsets. It transforms each quantitative item into fuzzy membership values and uses fuzzy operations to find fuzzy rules. The greater part of the algorithm that extracts association rules works in two phases: in the first one, candidate itemsets are generated and counted by scanning the transactions. If the number of an itemset appearing in the transactions is larger than a predefined threshold value (*minsup*), the itemset is thought as a large itemset. Itemsets with only one item are first processed. The large itemsets with one item are then combined to form candidate itemsets of two items. This process is repeated until all large itemsets are found. In the second phase, the desired association rules are induced from the large itemsets found in the first phase. All the possible combination ways of association rules for each large itemset are formed, and the ones with their calculated confidence values larger than a predefined threshold (*minconf*) are output as desired association rules.

The proposed solution for finding the association rules in terms of fuzzy terms from quantitative values is shown in Fig1. This approach consists of below phases.

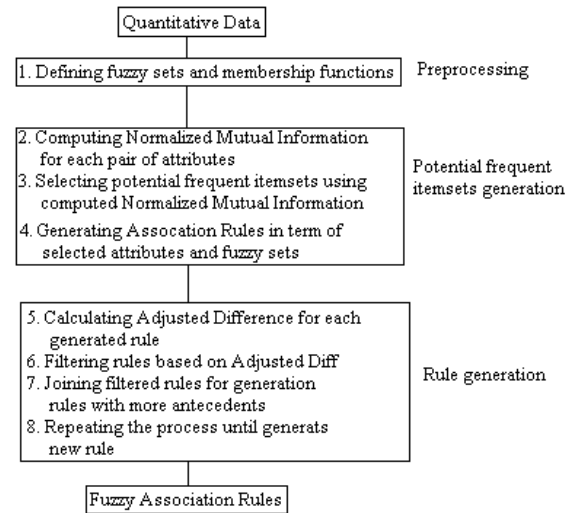


Fig1: Steps of proposed algorithm

A. Preprocessing using Fuzzy k-means clustering

Continuous attributes can be handled using fuzzy set theory. It is difficult to define the membership functions for each and every attribute based on intuition [10]. Hence, a clustering based approach (fuzzy k-means clustering) [6] has been used for finding membership for each attribute value. This method divides the values of each attribute into k -clusters [21]. The steps given below are used for clustering the attribute values:

- 1) Place k points into the space represented by the objects that are going to be clustered. These points represent initial group centroids.

- 2) Assign each object to the group that has the closest centroid.
- 3) When all objects have been assigned, recalculate the positions of the k-centroids.
- 4) Repeat steps 2 and 3 until the centroids move no longer. This produces a separation of the objects into groups from which the metric to be minimized can be calculated.
- 5) And this algorithm aims at minimizing the objective function (squared error function [21]).

$$J = \sum_{i=1}^N \sum_{j=1}^k (u_{ij})^m (x_i - c_j)^2 \quad (7)$$

where $(x_i - c_j)^2$ is a chosen distance measure between attribute value $x_i^{(j)}$ and the cluster center c_j , is an indicator of the distance of the N (number of records) attribute values from their respective cluster center. The resultant clusters have to be associated with k linguistic terms. These linguistic terms are associated based on cluster centers and attribute nature. Then the membership value is calculated for each value.

B. Compute Normalized Mutual Information

In this phase, the concepts of entropy and mutual information were applied that originates from information theory [20] in the context of QAR mining.

M is used as a measure to evaluate the strongness of the relationship between two attributes in a QAR mining problem. Given a predefined threshold μ , if $M \geq \mu$, we say that the two attributes are strongly related to each other; otherwise, we say that they are not strongly related. Ideally, M is a measure being able to identify attributes that do not constitute any significant QARs. Thus, we do not need to consider joining these attributes to produce candidate frequent itemsets in the mining process. Defining M as the mutual information between the attributes seems to be an ideal approach because mutual information, by definition, naturally measures the information that one attribute tells about another. For two attributes appearing in the same QAR, the strongness of their relationship is reflected by their mutual information.

However, as shown in [22] there are two crucial problems in the application of mutual information as such a measure of M. To tackle the mentioned problems, [22] has been proposed a normalization for mutual information (formula 6). Normalized mutual information gives the threshold μ an intuitive meaning and makes it relatively independent of specific attributes. Now the threshold μ indicates the minimum percentage of reduction in uncertainty about an attribute due to the knowledge of another attribute.

In this phase we compute values of normalized mutual information for all attribute pairs using formula 6.

C. Rule Generation and Rule selection

This phase deals with the generation and optimization of the rules. The combination of a pair of rules is under below conditions:

- 1) The consequent of two rules must be identical.
- 2) The rules must not contain similar antecedents on their left-hand sides.
- 3) The normalized mutual information of antecedents of two rules is greater than μ .

D. Identification of Interesting Associations

In order to decide whether the association between an attribute value, $r_{jk} \in \text{dom}(R_j)$, and another attribute value, $r_{pq} \in \text{dom}(R_p)$, is interesting, adjusted difference method [23] has been used. This is defined as:

$$d_{r_{pq}r_{jk}} = \frac{Z_{r_{pq}r_{jk}}}{\sqrt{\gamma_{r_{pq}r_{jk}}}} \quad (8)$$

$Z_{r_{pq}r_{jk}}$ is the *standardized difference* [23] given by

$$Z_{r_{pq}r_{jk}} = \frac{\text{count}_{r_{pq}r_{jk}} - e_{r_{pq}r_{jk}}}{\sqrt{e_{r_{pq}r_{jk}}}} \quad (9)$$

$e_{r_{pq}r_{jk}}$ is the number of records expected to have r_{jk} and r_{pq} calculated by

$$e_{r_{pq}r_{jk}} = \frac{\sum_{i=1}^m \text{count}_{r_{pq}r_{ji}} \sum_{i=1}^n \text{count}_{r_{pi}r_{jk}}}{M} \quad (10)$$

where m and n are number of linguistic terms defined on attributes j and p and $\gamma_{r_{pq}r_{jk}}$ is the *maximum likelihood estimate* [23] of the variance of $Z_{r_{pq}r_{jk}}$ and given by

$$\gamma_{r_{pq}r_{jk}} = \left(1 - \frac{\sum_{i=1}^m \text{count}_{r_{pq}r_{ji}}}{M}\right) \left(1 - \frac{\sum_{i=1}^n \text{count}_{r_{pi}r_{jk}}}{M}\right) \quad (11)$$

and $M = \sum_{u=1}^n \sum_{i=1}^m \text{count}_{r_{pu}r_{ji}}$

If $d_{r_{pq}r_{jk}} > 1.96$ (the 95 percentiles of the normal distribution) we can conclude that the association between r_{jk} and r_{pq} is interesting.

E. Algorithm in details

The proposed algorithm transforms each quantitative value into a fuzzy set with linguistic terms using membership functions, and then calculates the *Normalized Mutual Information* of each attribute on all the transaction data. Using these NMI prunes search space. The detail of the proposed mining algorithm is described as follows:

- 1) Computation all the values of normalized mutual information between each distinct pair of attributes.
- 2) If x_i, x_j are two adjusted attributes then $\tilde{I}(x_i, x_j) \geq \mu$ represents the strong information relationship between the attributes in a QAR mining problem. We provide the user with the flexibility to specify the threshold μ to be a

value in the range of $[0, 1]$, according to the user's requirement of the strongness of the relationship between the attributes. Relation $\tilde{l}(x_i, x_j) \geq \mu$ represents the set of attributes which are potential frequent itemsets. Essentially we utilize *mutual information* to do the pruning at the attribute level. Only the attribute set which have MI greater than μ are considered to generate rule. Meanwhile, we also check the support condition of the itemsets to make sure that they are frequent.

- 3) Given MI list, we construct rules level by level as follows:
 - a) After determining potential frequent itemsets, a set of rules with one antecedent is generated. The identification of interesting association is based on an objective interesting measure called *adjusted deference* [23]. This measure is employed to determine whether the association between linguistic term R_{pq} and another linguistic term R_{jk} is interesting. Combining those rules in the next step, results in the set of 2-dimensional candidate rules.
 - b) Any possible combination of the 1 and 2-antecedent rules that have the same consequent and not containing common antecedents would be a candidate rule with 3-antecedents. Only if all 2-dimensional sub-rules of a rule are present in the candidate set of the previous stage, a rule will be evaluated. The normalized mutual information is used to avoid the time-consuming evaluation of some useless rules.
 - c) Similarly, generating a rule with n-antecedent is performed by the combination of candidate rules with n-1 and 1-antecedent rules.

F. Experimental evaluation

We evaluate the performance of our algorithm on real datasets. We use Mining Fuzzy Weighted Association Rules (MFAR) [18] algorithm as the baseline for comparison on the efficiency of the algorithms. Real datasets are chosen from the commonly used UCI machine learning repository [24]. Table 1 lists the name, the number of attributes and the number of transactions of all datasets. Also, we use Loan data (introduced in [18]) for comparing quality of the mined rules. All the experiments are run on an XP machine with a 1.8 GHz Intel and 1GB RAM.

Table 1: Dataset Description

Dataset	No. of transaction	No. of attributes(QA)
Letter-recognition	20000	17(16)
Yeast	1484	9(8)
Loan	650	5(5)

In proposed algorithm we don't need *minsup* and *minconf* measures for evaluation the rules, but for comparing runtime of algorithms, we generate various sets of QARs at the minimum confidence and minimum support thresholds. The number of association rules decreases along with an increase in *minsup* (or *minconf*) under a given specific *minconf* (or *minsup*), which shows an appropriate *minsup* (or *minconf*) can constraint the number of association rules and avoid the

occurrence of some association rules so that it cannot yield a decision. These results have shown in figures 2-3. The results are as expected and quite consistent with our intuition.

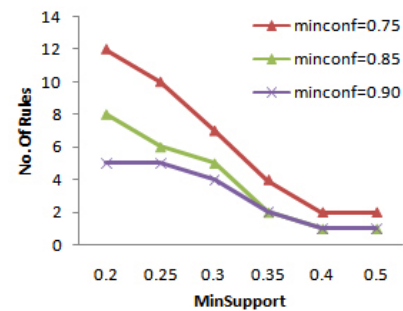


Fig.2. Association rules and *Minsup* in MFAMI algorithm

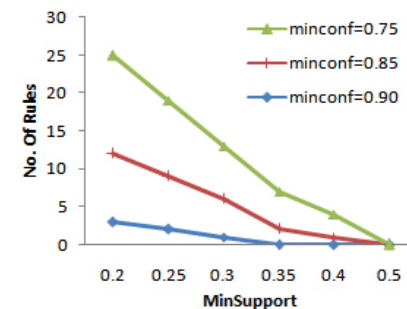


Fig.3. Association rules and *Minsup* in MFA algorithm

Figures 4-6 show the running time (generating frequent itemsets and finding rules) for each dataset.

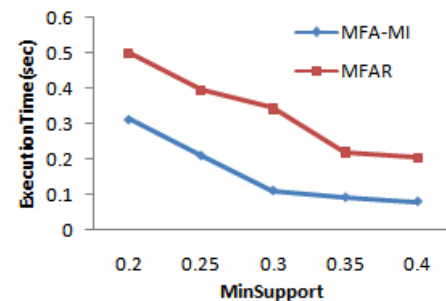


Fig. 4. Performance on Loan Dataset

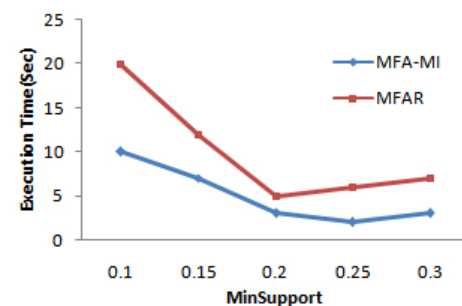


Fig.5. Performance on Letter-recognition dataset

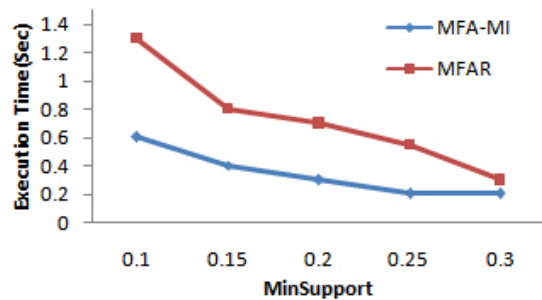


Fig.6. Performance on Yeast dataset

V. CONCLUSION

In this paper by using mutual information and fuzzy theory, we propose an efficient algorithm for discovery quantitative association rules. As said in [22], we apply MI to discover the informative relationship between the attributes in a QAR mining problem. The mutual information enumeration limits the mining process to a smaller but more relevant search space. Using the proposed method for rule generation, will be possible to generate rules without frequent itemset generation. By using potential frequent itemsets, we generate 1-antecedent rules. Through joining 1-dimensional rules with rules of each step, rules with more antecedents would be generated. Also we have used *adjusted difference* for finding the interestingness among the attributes in order to generate the rules.

REFERENCES

- [1] R. Agrawal and R. Smkant, "Fast algorithms for mining association rules", In Proceedings of the 20th International Conference on Very Large Databases, Santiago, Chile, September 1994, pp. 487-499.
- [2] J. S. Park and M. -S. Chen and P.S. Yu, "An effective hash-based algorithm for mining association rules", in proceedings of the ACM SIGMOD International Conference on Management of Data", San Jose, CA, May 1995, pp 175-186.
- [3] A Savasere, E. Ommicinski and S Navathe, "An efficient algorithm for mining association rules in large databases", In Proceedings Of the 21st International Conference on Very Large Databases, Zurich, Switzerland, September 1995, pp 432-444.
- [4] Agrawal R, Imielinski T, Swami A, "Mining association rules between sets of items in large databases". Buneman P, Jajodia S (eds) Proceedings of the ACM SIGMOD international conference on management of data, Washington DC, May 1993, pp 207-216.
- [5] Srikant R, Agrawal R, "Mining quantitative association rules in large relational tables". In: Proceedings of the ACM SIGMOD international conference, Canada, 1996, pp 1-12.
- [6] Ng RT, Han J, "Efficient and effective clustering methods for spatial data mining". In: Proceedings of the 20th VLDB conference, 1994, pp 144-155.
- [7] Gyenesei A, "Mining weighted association rules for fuzzy quantitative items", TUCS Technical Report No. 346, 2000, pp 1-12
- [8] P. Bosc and O. Pivert, "On some fuzzy extensions of association rules", Proceedings of IFSA-NAFIPS 2001, Piscataway, NJ, IEEE Press, 2001, pp. 1104-9.
- [9] T.-P. Hong, C.-S. Kuo and S.-C. Chi, "Trade-off between computation time and numbers of rules for fuzzy mining from quantitative data", International Journal of Uncertainty, Fuzziness and Knowledge-based Systems 9, 2001 587-604.
- [10] C. Kuok, A. Fu and H. Wong, "Mining fuzzy association rules in databases", ACM SIGMOD Record, 27, 1998, pp. 41- 46.
- [11] R. Ladner, F.E. Petry and M.A. Cobb, "Fuzzy set approaches to spatial data mining of association rules", Transactions in GIS, 7, 2003, pp. 123-138.
- [12] J. Shu, E. Tsang and D. Yeung, "Query fuzzy association rules in relational databases", In Proceedings of IFSANAFIPS 2001 Piscataway, NJ, IEEE Press, 2001, pp. 2989- 2993.
- [13] W. Zhang, "Mining fuzzy quantitative association rules", Proceedings of IEEE International Conference on Tools with Artificial Intelligence 1999 Piscataway, NJ, IEEE Press, 1999, pp. 99-102.

- [14] Chan KCC, Au W, "Mining fuzzy association rules". In: Proceedings of the sixth international conference on information and knowledge management, 1997, pp 209-215.
- [15] Yager RR, "On linguistic summaries of data, knowledge discovery in database", 1996, pp 347-363.
- [16] Ishibuchi H, Nakashima T, Murata T, "Three-objective genetics-based machine learning for linguistic rule extraction", 2001, Inf Sci 136:109-133.
- [17] Kaya M, Alhajj R, "Facilitating fuzzy association rules mining by using multi-objective genetic algorithms for automated clustering", In: Proceedings of the third IEEE international conference on data mining (ICDM'03), 2003, pp 561-564.
- [18] D.L. Olson, Yanhong Li, "Mining Fuzzy Weighted Association Rules", Proceedings of the 40th Hawaii International Conference on System Sciences, 2007.
- [19] T.P. Hong, C.S. Kuo, S.C. Chi, "Mining association rules from quantitative data", Intell. Data Anal. 3 (5) 1999, 363-376.
- [20] Shannon C, "A mathematical theory of communication, i and ii". Bell Syst Tech J 27:379-423, 623-656, July, October 1948.
- [21] Bezdek JC, "Pattern recognition with fuzzy objective function algorithms", Plenum, New York 1981.
- [22] Yiping Ke · James Cheng · Wilfred Ng, "An information-theoretic approach to quantitative association rule mining", 22 July 2007, Springer-Verlag London Limited
- [23] Au W-H, Chan KCC, "FARM: a data mining system for discovering fuzzy association rules", In: Proceedings of 8th IEEE international conference on fuzzy systems, Seoul, Korea 1999, pp 1217-1222
- [24] Asuncion A, Newman DJ, "UCI machine learning repository", Irvine, CA University of California, Department of Information and Computer Science, 2007, <http://www.ics.uci.edu/~mllearn/MLRepository.html>.