



Use of Vague Set Theory for Multiple Objectives Decision Making

Rupali Rajput, Vivek jain

Department of CSE and IT, SRCEM
Gwalior (M.P.), India

Abstract- In this paper, we aim to extend the notion of classical soft expert sets to possibility vague soft expert sets by applying the theory of soft expert sets to possibility vague soft sets. The complement, union, intersection, AND and OR operations as well as some related concepts pertaining to this notion are defined. Since vague sets can provide more information than fuzzy sets, it is superior in mathematical analysis of system with uncertainty. In this paper, a new method vague set-based is proposed to deal with decision fusion problem. Lastly, this concept is applied to a decision making problem and its effectiveness is demonstrated using a hypothetical example.

General Terms: Data Mining, Vague Sets, Vague Association Rule Mining, Multiple Objectives Decision Making, Objective Decision Making Matrix.

I. INTRODUCTION

A. Data Mining

Data mining refers to extracting or “mining” knowledge from huge volume of data [1]. By performing information mining, interesting knowledge, reliabilities, or high-level data can be extracted from database and viewed or browsed from various method. Data mining is regarded as one of the essential frontiers in database framework and one of the encouraging integrative evolution in the information industry.

Rule mining is a characteristic of data mining as well as a process of Knowledge Discovery in Databases (KDD), where various available data sources are explored, [2]. From this feature the concept of “Association Rule Mining” evolved. Association rule mining determines interesting association or correlation relationship among a large data set of items.

B. Fuzzy Sets

Fuzzy association rule mining primary started in the form of knowledge discovery in Fuzzy expert systems. A fuzzy expert system [3] uses a collection of fuzzy membership functions and rules, instead of Boolean logic, to reason about data [4]. The rules [5] in a fuzzy expert system are usually of a form similar to the following: “If it is raining then put up your umbrella” Here if part is the antecedent part and then part is the consequent part. This type of rules as a set helps in pointing towards any solution with in the solution set. Attribute values are not represented by just 0 or 1. Here attribute values are represented with in a range between 0 and 1.

There are many statuses of a piece of hesitation information (called *hesitation status(HS)*). Let us consider a motivating example of an online shopping scenario that involves various statuses: (s_1) HS of the items that the customer browsed only once and left; (s_2) HS of the items that are browsed in detail (e.g., the figures and all specifications) but not put into their online shopping carts; (s_3) HS of the items that customers put into carts and were checked out eventually. For example, given a criterion as the possibility that the customer buys an item, we have $s_1 \leq s_2 \leq s_3$.

The hesitation information can then be used to design and implement selling strategies that can potentially turn those “interesting” items into “under consideration” items and “under consideration” items into “sold” items.

Our modelling technique of HSs of an item rests on a solid foundation of vague set theory [6-8]. The main benefit of this approach is that the theory addresses the drawback of a single membership value in *fuzzy set theory* by using interval-based membership that captures three types of evidence with respect to an object in a universe of discourse: *support*, *against* and *hesitation*. Thus, we naturally model the hesitation information of an item in the mining context as the evidence of hesitation with respect to an item. To study the relationship between the support evidence and the hesitation evidence with respect to an item, we propose *attractiveness* and *hesitation* of an item, which are derived from the vague membership in vague sets. An item with high attractiveness means that the item is well sold and has a high possibility to be sold again next time.

Using the attractiveness and hesitation of items, we model a database with hesitation information as an *AH*-pair database that consists of *AH*-pair transactions, where *A* stands for attractiveness and *H* stands for hesitation. Based on the *AH*-pair database, we then discuss the notion of *Vague Association Rules (VARs)*, which capture four types of relationships between two sets of items: the implication of the attractiveness/hesitation of one set of items on the attractiveness/hesitation of the other set of items.

Vague Sets

Let I be a classical set of objects, called the universe of discourse, where an element of I is denoted by x .

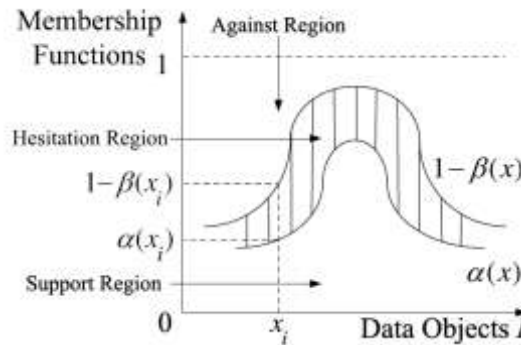


Fig 1. The true (α) and False (β) Membership Functions of a Vague set

(Vague Set) Gau's and Buehrer[9] introduced the notion of vague sets. A vague set V in a universe of discourse I is characterized by a true membership function, β_V , and a false membership function, α_V , as follows:

$\alpha_V: I \rightarrow [0, 1]$, $\beta_V: I \rightarrow [0, 1]$, where

$\alpha_V(x) + \beta_V(x) \leq 1$,

$\alpha_V(x)$ is a lowerbound on the grade of membership of x derived from the evidence for x , and $\beta_V(x)$ is a lower bound on the grade of membership of the negation of x derived from the evidence against x . Suppose $I = \{x_1, x_2, \dots, x_n\}$.

The grade of membership of x is bounded to $[\alpha_V(x); 1 - \beta_V(x)]$, which is a subinterval of $[0, 1]$ as depicted in Fig. 1. For brevity, we omit the subscript V from α_V and β_V . We say that to $[\alpha_V(x); 1 - \beta_V(x)]/x$ is a *vague element* and the interval to $[\alpha_V(x); 1 - \beta_V(x)]$ is the *vague value* of the object x .

C. Vague Association Rule Mining

In this section, the concept of Hesitation Statuses (HSs) of an item is shown and discussed how to model HSs. Then the notion of *Vague Association Rules* (VARs) and four types of support and confidence used in order to fully evaluate their quality. Some properties of VARs that are useful to improve the efficiency of mining VARs are presented.

Hesitation Information Modeling:-

A *Hesitation Status (HS)* is a specific state between two certain situations of "buying" and "not buying" in the process of a purchase transaction.

In order to capture the hesitation evidence and the hesitation order, a subinterval of $[\alpha_V(x); 1 - \beta_V(x)]$ is used to represent the customer's *intent* of each item with respect to different HSs. To obtain the intent value, the idea of linear extensions of a partial order is used.

Attractiveness and Overall Attractiveness:-

The attractiveness of x with respect to an HS_i , denoted as $att(x, s_i)$ is defined as the median membership of x with respect to S_i that is $\frac{1}{2} (\alpha_i(x) + (1 - \beta_i(x)))$.

The overall attractiveness of x is a function $ATT(x): I \rightarrow [0, 1]$ such that $ATT(x) = \frac{1}{2} (\alpha(x) + (1 - \beta(x)))$.

Given the intent $[\alpha_V(x); 1 - \beta_V(x)]$ of an item x for an $HS s_i$, we have a one-one corresponding pair of the attractiveness and hesitation of x , called the *AH-pair*, denoted as $[att(x; s_i); h(x)]$. Attractiveness and hesitation are two important concepts, since people may have special interest in finding ARs with items of high attractiveness (sold well) or high hesitation (almost sold).

Vague Association Rules and their Support and Confidence

We now present the notion of VARs and define the support and confidence of a VAR. Definition 6. (Vague Association Rule) A Vague Association Rule (VAR), $r = (X \rightarrow Y)$, is an association rule obtained from an AH-pair database.

Based on the attractiveness and hesitation of an item with respect to an HS, we can define different types of support and confidence of a VAR. We define *Attractiveness-Hesitation (AH)* support and *AH* confidence of a VAR to evaluate the VAR. Similarly, we can obtain the association between an itemset with high hesitation and another itemset with high attractiveness, between two itemsets with high attractiveness, and between two itemsets with high hesitation for different purposes.

Definition 3: (Hesitation and overall Hesitation)

Given an item $x \in I$ and a set of HSs $S = \{s_1, s_2, \dots, s_n\}$ with a partial order $\leq$. The hesitation of x with respect to a hesitation status $HS s_i \in S$ is a function $h_i(x): I \rightarrow [0, 1]$ such that $\alpha(x) + \beta(x) + \sum_{i=1}^n h_i(x) = 1$ where $h_i(x)$ represents the evidence for the HS s_i of x . The overall hesitation of x with respect to S is given by $H(x) = \sum_{i=1}^n h_i(x)$. This can be easily find from the above definition that $H(x) = 1 - \alpha(x) - \beta(x)$.

Definition 5: (Attractiveness and overall Attractiveness)

The attractiveness of x with respect to an HS s_i , denoted as $att(x, s_i)$ is defined as the median membership of x with respect to S_i that is $\frac{1}{2} (\alpha_i(x) + (1 - \beta_i(x)))$. The overall attractiveness of x is a function $ATT(x): I \rightarrow [0, 1]$ such that $ATT(x) = \frac{1}{2} (\alpha(x) + (1 - \beta(x)))$.

Definition 6: (Weighted attributes)

These are variables selected to calculate weight are known as weighting attributes $A(a_1, a_2, \dots, a_n)$ depending on domain it could be any variable such as item weight in case of supermarket domain.

Definition 7: (Item weight)

Item weight is the value attached to items representing its significance. In case of supermarket setting it can be the profit per unit sale of certain item. The item weight is function of selected weighting attributes. If item weight is $w(i)$ then $w(i) = f(a)$.

Definition 8: (Itemset weight)

Weight of an itemset is the weights of its enclosing items. It can be denoted as $w(is)$. The item weight is a special type of itemset weight when itemset has only one item. The average value of item weight is given by

$$w(is) = \frac{\sum_{k=1}^N w(i_k)}{N}$$

Definition 9: (Transaction weight)

Transaction weight is a type of itemset weight that is attached to each of the transactions. Higher transaction weight means more contribution in mining results. Considering scenario of supermarket the weight can be significance of a customer who made a certain transaction.

Definition 10: (Weighting spaces)

Items can be weighted within different weighting spaces depending on different scenario and mining focus, it is the context within which the weight are evaluated.

- *Transaction space* (WS_T) is defined for transactions rather than for items.
- *Item space* (WS_I) refers to space of the items collection that covers all the items appears in the transactions.
- *Inner – transaction space* (WS_t) is the space refers to host transaction that an item is weighted in.

Definition 11: (AH-pair transaction and database)

An AH-pair database is sequence of AH-pair transactions. An AH-pair transaction T is a tuple $\langle v_1, v_2, \dots, v_m \rangle$ on an itemset $I_T = \{x_1, x_2, \dots, x_m\}$ where $I_T \subseteq I$ and $v_j = \langle M_A(x_j), M_H(x_j) \rangle$ is an AH-pair of the item x_j with respect to a given HS or the overall hesitation for $1 \leq j \leq m$.

D. Multi-Attribute Decision Making

A General Overview Multi-Attribute Decision Making is the most well-known branch of decision making. It is a branch of a general class of Operations Research (or OR) models which deal with decision problems under the presence of a number of decision criteria. This super class of models is very often called multi-criteria decision making (or MCDM). According to many authors (see, for instance, [10]) MCDM is divided into Multi-Objective Decision Making (or MODM) and Multi-Attribute Decision Making (or MADM). MODM studies decision problems in which the decision space is continuous. A typical example is mathematical programming problems with multiple objective functions. The

first reference to this problem, also known as the "vector-maximum" problem, is attributed to [11]. On the other hand, MADM concentrates on problems with discrete decision spaces. In these problems the set of decision alternatives has been predetermined.

Although MADM methods may be widely diverse, many of them have certain aspects in common [12]. These are the notions of alternatives, and attributes (or criteria, goals) as described next.

Alternatives: Alternatives represent the different choices of action available to the decision maker. Usually, the set of alternatives is assumed to be finite, ranging from several to hundreds. They are supposed to be screened, prioritized and eventually ranked.

Multiple attributes: Each MADM problem is associated with multiple attributes. Attributes are also referred to as "goals" or "decision criteria". Attributes represent the different dimensions from which the alternatives can be viewed.

In cases in which the number of attributes is large (e.g., more than a few dozens), attributes may be arranged in a hierarchical manner. That is, some attributes may be major attributes. Each major attribute may be associated with several sub-attributes. Similarly, each sub-attribute may be associated with several sub-sub-attributes and so on. Although some MADM methods may explicitly consider a hierarchical structure in the attributes of a problem, most of them assume a single level of attributes (e.g., no hierarchical structure).

Conflict among attributes: Since different attributes represent different dimensions of the alternatives, they may conflict with each other. For instance cost may conflict with profit, etc.

Incommensurable units: Different attributes may be associated with different units of measure. For instance, in the case of buying a used car, the attributes "cost" and "mileage" may be measured in terms of dollars and thousands of miles, respectively. It is this nature of having to consider different units which makes MADM to be intrinsically hard to solve.

Decision weights: Most of the MADM methods require that the attributes be assigned weights of importance. Usually, these weights are normalized to add up to one.

Decision matrix: An MADM problem can be easily expressed in matrix format. A decision matrix A is an $(M \times N)$ matrix in which element a_{ij} indicates the performance of alternative A_i when it is evaluated in terms of decision criterion C_j , (for $i = 1, 2, 3, \dots, M$, and $j = 1, 2, 3, \dots, N$). It is also assumed that the decision maker has determined the weights of relative performance of the decision criteria (denoted as W_j , for $j = 1, 2, 3, \dots, N$). This information is best summarized in figure 1. Given the previous definitions, then the general MADM problem can be defined as follows [Zimmermann, 1991]:

Definition 1-1:

Let $A = \{ A_i, \text{ for } i = 1, 2, 3, \dots, M \}$ be a (finite) set of decision alternatives and $G = \{ g_i, \text{ for } j = 1, 2, 3, \dots, N \}$ a (finite) set of goals according to which the desirability of an action is judged. Determine the optimal alternative A^* with the highest degree of desirability with respect to all relevant goals g_i .

Criteria

	C_1	C_2	C_3		C_N
<u>Alt.</u>	W_1	W_2	W_3		W_N
A_1	a_{11}	a_{12}	a_{13}		a_{1N}
A_2	a_{21}	a_{22}	a_{23}		a_{2N}
A_3	a_{31}	a_{32}	a_{33}		a_{3N}
.	.	.	.	.	.
.	.	.	.	.	.
.	.	.	.	.	.
A_3	a_{M1}	a_{M2}	a_{M3}		a_{MN}

Figure 1: A Typical Decision Matrix.

Very often, however, in the literature the goals g_i are also called decision criteria, or just criteria (since the alternatives need to be judged (evaluated) in terms of these goals). Another equivalent term is attributes. Therefore, the terms MADM and MCDM have been used very often to mean the same class of models (i.e., MADM). For these reasons, in this paper we will use the terms MADM and MCDM to denote the same concept.

Multi-attribute decision making methods

Consider a multi-attribute decision making problem with m criteria and n alternatives. Let $C_1, \dots, C_m$ and $A_1, \dots, A_n$ denote the criteria and alternatives, respectively. A standard feature of multi-attribute decision making methodology is the *decision table* as shown below. In the table each row belongs to a criterion and each column describes the performance of an alternative. The score a_{ij} describes the performance of alternative A_j against criterion C_i . For the sake of simplicity we assume that a higher score value means a better performance since any goal of minimization can be easily transformed into a goal of maximization.

The following eleven MCDM methods were identified throughout the review: 1) Multi-Attribute Utility Theory, 2) Analytic Hierarchy Process, 3) Fuzzy Set Theory, 4) Case-based Reasoning, 5) Data Envelopment Analysis, 6) Simple Multi-Attribute Rating Technique, 7) Goal Programming, 8) ELECTRE, 9) PROMETHEE, 10) Simple Additive Weighting, and 11) Technique for Order of Preference by Similarity to Ideal Solution. The following sections address each particular method first with a summary and discussion of the reviewed studies, and then follow with a brief discussion of the general approach and an examination of the advantages and disadvantages of each method.

E The Analytic Hierarchy Process

The Analytic Hierarchy Process (AHP) was proposed by Saaty [13]. The basic idea of the approach is to convert subjective assessments of relative importance to a set of overall scores or weights. AHP is one of the more widely applied multiattribute decision making methods.

Features of the AHP

The AHP is a very flexible and powerful tool because the scores, and therefore the final ranking, are obtained on the basis of the pairwise relative evaluations of both the criteria and the options provided by the user. The computations made by the AHP are always guided by the decision maker's experience, and the AHP can thus be considered as a tool that is able to translate the evaluations (both qualitative and quantitative) made by the decision maker into a multicriteria ranking. In addition, the AHP is simple because there is no need of building a complex expert system with the decision maker's knowledge embedded in it. On the other hand, the AHP may require a large number of evaluations by the user. Although every single evaluation is very simple, since it only requires the decision maker to express how two options or criteria compare to each other, the load of the evaluation task may become unreasonable. In fact the number of pairwise comparisons grows quadratically with the number of criteria and options. For instance, when comparing 10 alternatives on 4 criteria, $4 \cdot 3/2 = 6$ comparisons are requested to build the weight vector, and $4 \cdot (10 \cdot 9/2) = 180$ pairwise comparisons are needed to build the score matrix. However, in order to reduce the decision maker's workload the AHP can be completely or partially automated.

Implementation of the AHP

The AHP can be implemented in three simple consecutive steps:

- 1) Computing the vector of criteria weights.
- 2) Computing the matrix of option scores.
- 3) Ranking the options.

Consider how to derive the weights of the criteria. Assume first that the m criteria are not arranged in a tree-structure. For each pair of criteria, the decision maker is required to respond to a pairwise comparison question asking the relative importance of the two. The responses can use the following nine-point scale expressing the intensity of the preference for one criterion versus another

Value of c_{ij}	Interpretation
1	i and j are equally important
3	i is slightly more important than j
5	i is more important than j
7	i is strongly more important than j
9	i is absolutely more important than j

Table 1. Table of relative scores

If the judgement is that criterion C_j is more important than criterion C_i , then the reciprocal of the relevant index value is assigned.

Let c_{ij} denote the value obtained by comparing criterion C_i relative to criterion C_j . Because the decision maker is assumed to be consistent in making judgements about any one pair of criteria and since all criteria will always rank equally when compared to themselves, we have $c_{ij}=1/c_{ji}$ and $c_{ii}=1$.

This means that it is only necessary to make $1/2m(m-1)$ comparisons to establish the full set of pairwise judgements for m criteria. The entries c_{ij} , $i, j=1, \dots, m$ can be arranged in a *pairwise comparison matrix* C of size $m \times m$.

The next step is to estimate the set of weights that are most consistent with the relativities expressed in the comparison matrix. Note that while there is complete consistency in the (reciprocal) judgements made about any one pair, consistency of judgements between pairs, i.e. $c_{ij}c_{jk}=c_{ik}$ for all i, j, k , is not guaranteed. Thus the task is to search for an m -vector of the weights such that the $m \times m$ matrix W of entries w_i/w_j will provide the best fit to the judgments recorded in the pairwise comparison matrix C . Several of techniques were proposed for this purpose.

Saaty's original method to compute the weights is based on matrix algebra and determines them as the elements in the eigenvector associated with the maximum eigenvalue of the matrix. The eigenvalue method has been criticized both from prioritization and consistency points of view and several other techniques have been developed. A number of other methods are based on the minimization of the distance between matrices C and W . Some of these approaches give the vector w directly or by simple computations, some other ones require the solution of numerically difficult optimization problems. One of these approaches, the logarithmic least squares method, results in a straightforward way of computing vector w : calculate the geometric mean of each row in the matrix C , calculate the sum of the geometric means, and normalize each of the geometric means by dividing by the sum just computed [14]. [15] references on distance-minimizing methods and a new approach based on singular value decomposition.

In the practice the criteria are often arranged in a tree-structure. Then, AHP performs a series of pairwise comparisons within smaller segments of tree and then between sections at a higher level in the tree-structure.

Similarly to calculation of the weights for the criteria, AHP also uses the technique based on pairwise comparisons to determine the relative performance scores of the decision table for each of the alternatives on each subjective (judgemental) criterion. Now, the pairwise questions to be answered ask about the relative importance of the performances of pairs of alternatives relating the considered criterion. Responses use the same set of nine index assessments as before, and the same techniques can be used as at computing the weights of criteria.

With the weights and performance scores determined by the pairwise comparison technique above, and after further possible normalization, alternatives are evaluated using any of the decision table aggregation techniques of the MAUT methods. The so-called additive AHP uses the same weighted algebraic means as SMART, and the multiplicative AHP is essentially based on the computation of the weighted geometric means.

A number of specialists have voiced a number of concerns about the AHP, including the potential internal inconsistency and the questionable theoretical foundation of the rigid 1-9 scale, as well as the phenomenon of rank reversal possibly arising when a new alternative is introduced. On the same time, there have also been attempts to derive similar methods that retain the strengths of AHP while avoiding some of the criticisms. See Triantaphyllou, E. [16] for state-of-art surveys and further references.

II. RELATED WORK

Bala Yesu Chilakalapudi, Narayana Satyala and Satyanarayana Menda [17] presented an algorithm for a resolving the issue problem of extracting frequent item sets from a huge vague database, interpreted under the Possible World Semantics (PWS). This issue is strictly difficult since a vague database consists of an exponential number of possible worlds. By examining the mining process can be modeled as a Poisson binomial distribution, an algorithm is implemented which can effectively and exactly determine frequent item sets in a huge vague database. The devised mining algorithm facilitate Probabilistic Frequent Item set (PFI) outcomes to be re-energized. The devised algorithm can maintain incremental mining and provides the precise outcomes on mining the vague database. The broad estimation on real data set to certify the scheme is performed.

Starr and Zeleny [18] give a brief outline of the origins of MODM in the field of management science. This work began in the early 1950's. After initial contributions, the next major contribution was that of goal programming. In goal programming, a multiplicity of objectives are reduced to a single objective by minimizing deviations of each objective from certain pre-specified target levels or goals. The following decade saw traditional utility extended to multiattribute utility theory, and with Johnsen's (1968) study on the multigoal nature of the firm, Starr and Zeleny (19) suggest that "multiple criteria decision making was firmly on its path."

In the naive approach used by Wallenius, the OM chooses a desired solution and is told only whether or not it is feasible; no attempt is made to find an efficient solution. And Martinson has used a minmax formulation where the objectives were normalized using the fractional achievement norm. In his solution method the OM was required to provide a set of weights which reflected the relative importance of each objective. These were then used to find the achieved solution. Some practical experience with this approach has indicated that the OM often has difficulty in relating the achieved solution to the particular set of weights chosen.

Considerable research effort has been directed to finding solution methods which ensure that only efficient solutions are generated; in fact in many MODM solution methods, the actual optimization involves nothing more than distinguishing between efficient and inefficient solutions. As will be seen from the literature review to follow, almost all MODM solution methods only consider efficient solutions; consequently the characterization of efficient solutions is of high priority. Kuhn and Tucker, in presenting necessary and sufficient conditions for solving the single objective optimization problem, also extended their work to the multiple objective case. Let $\lambda_i, i = 1, 2, \dots, m$ be the Lagrange multipliers for each constraint of X and assume that the objective functions are concave and the feasible set X is convex.

An Lu and Wilfred Ng [20] devised an algorithm for the issue, given a vague relation r over a schema R and a set of FDs F over R , what is the "best" approximation of r with respect to F when taking into account of the median membership (m) and the imprecision membership (i) thresholds. Employing these two thresholds of a vague set, defined the notion of mi-overlap among vague sets and a merge operation on r . Satisfaction of an FD in r is defined in terms of values being mi-overlapping. The main outcomes is that the output of the process is the most object-precise approximation of r with respect to F .

III. PROPOSED WORK

We have proposed an algorithm which uses MODM using the concepts of vague sets so as to optimize the results. The whole proposed process is explained here with the help of a flowchart and pseudocode that represents the flow of the proposed algorithm.

The proposed algorithm is explained below with the help of a flow chart. Also a pseudo code has been written for the algorithm designed. The flowchart can be shown as below:-

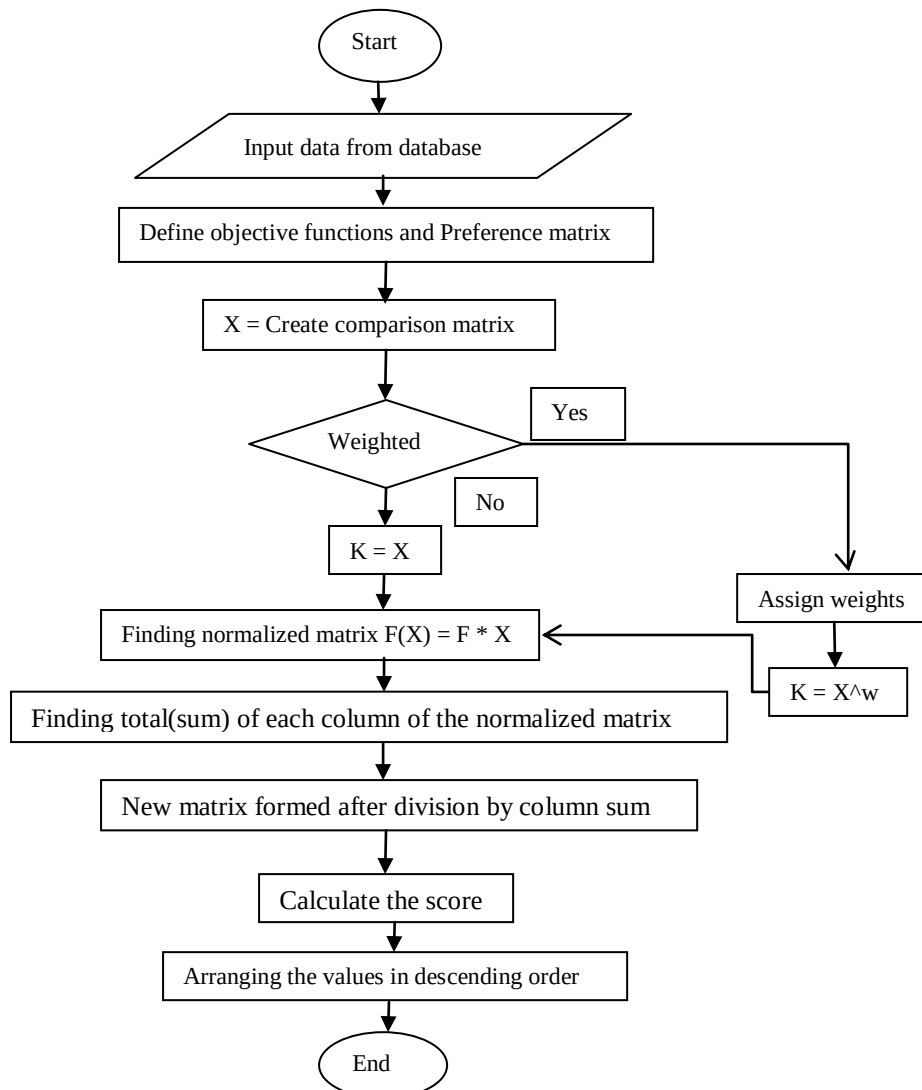


Fig: Flowchart of Proposed Algorithm

This flowchart explains the whole process of the proposed algorithm. New improved algorithm proposes the implementation of MODM to reduce the number of unnecessary itemsets. The pseudocode for this algorithm has been explained below:-

1. Algorithms for Mining vague association rules

Here we proposed an algorithm that mine the vague association rules from the given sample database in table. Then we apply the weighted concept to find the proficient rule that is used to increase the profitability concern of store. Firstly we mine set of all the values in the preference matrix.

Algorithm 1: MODM()

1. Calculate the F using objective functions.
2. Setting preference matrix
3. Calculating X i.e. pairwise comparison matrix.
4. Checking whether the matrix is weighted or unweighted.
5. For weighted $K=X$
6. $F(X) = F * K$
7. Finding total of each column of normalized matrix.
8. Matrix V = dividing each column by the sum.
9. Call vague().
10. Arrange in descending order based on scores.
11. End for.
12. End.

The second module of framework is to calculate the vague values that finds out the score so as to arrange the values according to the preferences. This algorithm takes intent as input which is calculated in previous algorithm.

Algorithm 2: Vague()

1. Setting the lower and upper bounds.
2. Calculating the values that are in favor, against and neutral. i.e. F, A and N values.
3. Calculating the truth and false value.
 - a. $\text{Truth}(t) = \text{Favour} / \text{total values}$
 - b. $\text{False value } (f) = \text{Against} / \text{total values}.$
 - c. $\text{Score} = t - f$
4. Arrange the values on the basis of score.
5. End.

IV. RESULT ANALYSIS

We have done experiments on base and proposed algorithms.

Simulation tool: - MATLAB R2013.

The dataset contains data for 1 year based on the requirement of the algorithm. The database with all the values is stored in Microsoft Office Excel 2013. All experiments were performed well and fully on Dell workstation with 4 GB RAM and 32-bit operating system, running windows 7.

V. RESULTS OBTAINED

Steps in the process of the implementation:-

Step 1: When clicking on run, a GUI is displayed which is shown below:-



Fig 2. GUI

Step 2: On clicking the unweighted button, the algorithm for unweighted items work behind the algorithm runs and gets its database from the excel sheet where database has been created. The patterns and results are generated which are shown below:-

Unweighted :

$x_4 > x_2 > x_1 > x_5 > x_3$

$1.000 > 1.000 > 0.675 > 0.652 > 0.389$

Weighted :

$x_4 > x_2 > x_5 > x_1 > x_3$

$1.000 > 1.000 > 0.475 > 0.421 > 0.398$

VI. CONCLUSION

This paper proposes a brief idea of vague set based approach for extracting patterns based on attractiveness and hesitation. Due to the use of this concept, an elaboration of the whole mining and the sub-fields of mining have been explained. The vague set mining is a new concept and a new proposed have been proposed on the basis of it. The optimization algorithm can has optimized the algorithm by giving more patterns i.e. converting the hesitation of higher order into attractiveness and thus giving user an idea about the possibility of the patterns to be generated.

REFERENCES

- [1] J. Han and M. Kamber , “Data Mining: Concepts and Techniques”, 2nd ed.,The Morgan Kaufmann Series in Data Management Systems, Jim Gray, Series Editor 2006.
- [2] Frawley, William J.; Piatetsky-Shapiro, Gregory; Matheus, Christopher J.: Knowledge Discovery in Databases: an Overview. AAAI/MIT Press, 1992.
- [3] Türksen, I.B. and Tian Y. 1993. Combination of rules and their consequences in fuzzy expert systems, Fuzzy Sets and Systems, No. 58,3-40, 1993.
- [4]<http://www.cs.cmu.edu/Groups/AI/html/faqs/ai/fuzzy/part1/faq-doc-4.html>
- [5] Gau, W.L., Buehrer, D.J.: Vague sets. IEEE Transactions on Systems, Man, and Cybernetics 23 (1993) 610–614
- [6] Lu, A., Ng, W.: Managing merged data by vague functional dependencies. In: ER. (2004) 259–272
- [7] Lu, A., Ng, W.: Vague sets or intuitionistic fuzzy sets for handling vague data: Which one is better In: ER. (2005) 401–416
- [8] W. L. Gau and D. J. Buehrer. Vague sets. IEEE Transactions on Systems, Man and Cybernetics, 23:610–614, 1993.

- [9] Pradnya A. Shirsath, Vijay Kumar Verma, "A Recent Survey on Incremental Temporal Association Rule Mining", International Journal of Innovative Technology and Exploring Engineering (IJITEE) ISSN: 2278-3075, Volume-3, Issue-1, June 2013.
- [10] Zimmermann, H.-J., *Fuzzy Set Theory and Its Applications*, Kluwer Academic Publishers, Second Edition, Boston, MA, 1991.
- [11] Kuhn, H.W. and A.W. Tucker, "Nonlinear Programming," *Proc. 2nd Berkeley Symp. Math. Stat. Prob.*, 481-492, 1951.
- [12] Chen, S.J. and C.L. Hwang, *Fuzzy Multiple Attribute Decision Making: Methods and Applications*, Lecture Notes in Economics and Mathematical Systems, No. 375, Springer-Verlag, Berlin, Germany, 1992.
- [13] Saaty, T.L., *The Analytic Hierarchy Process*, McGraw-Hill International, New York, NY, 1980.
- [14] Saaty, T.L. and Vargas, L.G. (1984) .Comparison of eigenvalue, logarithmic least squares and least squares methods in estimating ratios., *Mathematical Modelling*, 5, 309-324.
- [15] Gass, S. I. and Rapcsák, T. (2004) .Singular value decomposition in AHP., *European Journal of Operational Research*, 154, 573-584.
- [16] Triantaphyllou, E. (2000) *Multi-Criteria Decision Making Methods: A Comparative Study*, Kluwer Academic Publishers, Dordrecht.
- [17] Figueira, J., Greco, S. and Ehrgott, M. (Eds.) (2004) *Multiple Criteria Decision Analysis: State of the Art Surveys*, Springer, New York.
- [18] BalaYesuChilakalapudi, NarayanaSatyala and SatyanarayanaMenda, "An Improved Algorithm for Efficient Mining of Frequent Item Sets on Large Uncertain Databases", *International Journal of Computer Applications (0975 – 8887) Voluyume 73– No.12, July 2013*.
- [19] Miller, D.W., and M.K. Starr, *Executive Decisions and Operations Research*, Prentice-Hall, Inc., Englewood Cliffs, NJ, 1969.
- [20] An Lu and Wilfred Ng, "Vague Sets or Intuitionistic Fuzzy Sets for Handling Vague Data- Which One Is Better?" 2005 © Springer.