

# EXTRACTING PATTERN FROM LARGE DATA CORPUS USING APPROXIMATE REASONING AND INTUITIONISTIC FUZZY CLUSTERING TECHNIQUES

**Dr. ASHIT KUMAR DUTTA**

Associate Professor, Department of Computer Science, Shaqra University, Kingdom of Saudi Arabia

E-mail: drashitkumar@yahoo.com

## ABSTRACT

The fuzzy logic is a familiar method to automate a complex activity. Intuitionistic fuzzy clustering is the extension of fuzzy logic. Approximate reasoning is a concept to deal vague and complex data. Many real time problems were solved with the help of Intuitionistic fuzzy concepts. Classification and clustering are the well known methods in the process of extraction of knowledge from the large data corpus. The existing techniques that are based on clustering methods could not provide optimum results with less computation cost. The efficiency of the existing methods was not up to the mark on large datasets. The objective of the proposed research is to provide an efficient technique for the extraction of meaningful pattern from large dataset with least computation cost. The proposed method has used the Intuitionistic fuzzy clustering technique with approximate reasoning for the extraction of pattern from the benchmark datasets. The experiment and results on large data corpus has proved that the level of the proposed research is satisfactory.

**Keywords:** *Fuzzy clustering, Pattern extraction, Pattern Miner, Fuzzy – C Means, Soft computing*

## 1. INTRODUCTION

Classification, clustering, regression analysis, and association rules are the familiar algorithms available in the field of data mining. The clustering technique is the process of combining the group of objects into homogenous sets. The objects which are belonging to sets are more similar than the objects of other sets[1][2]. Partitioning, Hierarchical, Fuzzy, Density – based, and Model based clustering are the types of clustering methods. Fuzzy clustering (FC) and Model based clustering are the advanced clustering methods for the generation of pattern from the large dataset[3][4]. The fuzzy clustering will allow an object to be part of more than one cluster[5][6]. The individual object has a group of membership coefficient related to the level of cluster. The centroid of a cluster will be calculated by the weightage of clusters and mean of all points[7][8]. The fuzzy clustering algorithms are the alternatives of K – mean clustering algorithm.

Existing clustering algorithms are using the bag of words approach to extract pattern from large dataset[9][10]. Lexical chain semantic approaches are used to cluster the documents. It is used to

uncover the complexity over the documents and generate high dimensionality of the clusters.

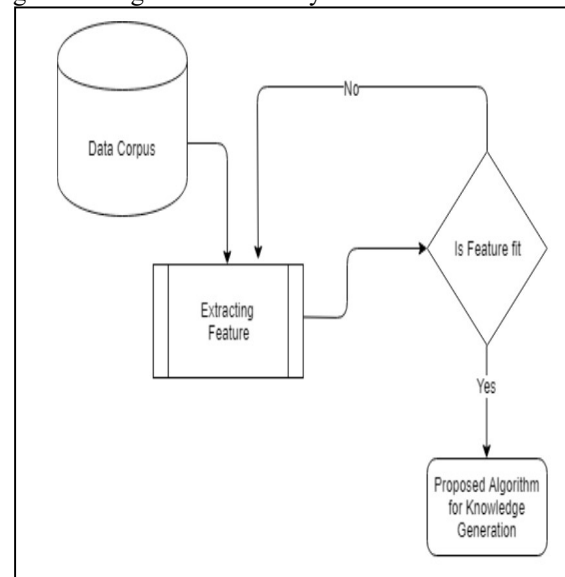


Figure 1 Data Extraction Process

The figure 1 represents the process of extraction of pattern from data corpus. The extraction of the features from the dataset is the important process in pattern extraction. Approximate Reasoning(AR) is used to filter the features extracted from the dataset.

The processes will be explained in section 3. The proposed algorithm will extract pattern from the selected features. The training phase will be used to teach the methods to predict interesting patterns.

The concept of Intuitionistic Fuzzy Clustering (IFC) is the extension of both fuzzy and classical sets[11][12][13]. Inclusion, Equality, Classical negation, Conjunction, and Disjunction are the operators and relations of the fuzzy sets which can be implemented on IFC. Atanassov and Prof. Krassmir.T were introduced IF Logic to study the properties of an object[14][15][16]. The IFS will use a hesitation margin rather than assuming negation value to the degree of non – membership objects. The theory of IFC will be useful when a linguistic variable of membership functions are vague and rough. It will be useful tool to handle uncertainty on objects[17][18][19]. Customer behaviour analysis, product promotion, financial management, medical imaging system, and psychological investigations are the real time applications based on IFC[20][21][22].

The aim of the research is to provide an efficient method to extract pattern from large data corpus. The research is organized as follows: Section – 2 will provide the review of literature. Section – 3 will give details about the methodology of IFC. Section – 4 will provide the data preparation and pre – process activities. Section – 5 will provide details about results and experiments on the benchmark datasets. Finally, the conclusion part will discuss the future implementation of the research.

## 2. REVIEW OF LITERATURE

A systematic approach was followed in the collection of literatures on IFC and pattern extraction. Google scholar, IEEE explore, and Springer are the portals used to search the research articles. “Intuitionistic Fuzzy Clustering”, “Approximate Reasoning”, “Pattern extraction”, and “Knowledge extraction” are the keywords used in the process of collection of research articles. The following table 1 will show the details about the research articles used in the research.

The acronym “GS” is used for Google Scholar, “IEEE” is for IEEE Explore, “SPR” is for Springer.

Table 1 Details of Research

| Years   | GS | IEEE | SPR | Citation |
|---------|----|------|-----|----------|
| 2013-14 | 3  | 2    | 2   | >50      |
| 2014-15 | 3  | 3    | 2   | >30      |
| 2015-16 | 2  | 1    | 3   | >30      |
| 2016-17 | 1  | 0    | 1   | >20      |
| 2017-18 | 2  | 1    | 1   | >15      |

The initial phase of literature collection has collected a total of 55 research articles. The following conditions were followed for the selection of research articles.

The following conditions were followed to filter the research articles.

- The studies that are published between 2013 and 2014 should have more than 50 citations.
- The studies that are published between 2014 and 2015 should have more than 30 citations.
- The studies that are published between 2015 and 2016 should have more than 30 citations.
- The studies that are published between 2016 and 2017 should have more than 20 citations.
- The studies that are published between 2017 and 2018 should have more than 15 citations.

By following the above criteria, a number of 18 articles were rejected from the total number of articles. Each article were comprehensively scanned and matched with objective of the proposed research. During this process, a number of 10 articles were rejected and arrived with a final number of 27 articles. The following details were extracted from the collected articles.

Peerasak Intarapaiboon[6] has proposed a method to classify text using similarity measures using IF sets. Membership and Non- membership are the complementary degrees of fuzzy sets. The IF sets will not agree with the fuzzy set on its nature of membership and non – membership functions. Text categorization is the concept of assigning a document to a set of categories to extract pattern. The author had developed a framework to classify text using IF sets to generate a pattern for each object in a set. The framework had training and testing phase for the process of extraction of pattern from dataset. The pre – processing of datasets were carried out in both training and testing phases. The vector and IF sets based representation were used for the pre – process of datasets. The pattern learning module will provide training for the system to generate patterns. In the testing phase, the pre – process unit will clean the dataset and directly classify the dataset. The drawback of the research is the computation time and accuracy. The pre – processing unit should not be involved in the pattern generation. It should be a separate entity. The accuracy of the testing phase will be less as pre – processing unit is involved in the phases of framework.

S.V. Aruna kumar and B.S.Harish[7] have proposed an algorithm for medical image segmentation. The Sugena’s and Yuger’s IF set generators are not satisfying the basic condition of

intuitionism. Medical imaging is a crucial activity to understand and analyse the human organs. X – ray, Computer Tomography, and Magnetic Resonance Imaging are the different tools available to obtain medical images. Segmentation is the process of identifying abnormal changes in tissues and organs. Thresholding, region growing, and merging applied for medical image segmentation are the classical image processing techniques. Vagueness in imprecise grey levels and object boundary are the samples of uncertainties involved in medical images. Classical fuzzy clustering will not handle the hesitation in medical images. The method which is used in the research had two steps: Intuitionistic representation of image and IFC method. The objective function of modified IFC method has used the method to cluster feature vector by computing local minima value. The modified hausdorff distance is used to evaluate the performance of methods. The experiment and results have shown the performance of the method was satisfactory.

Kuo – Ping Lin et.al[8]. have proposed a novel IF rough set model to enhance the rule generation. The rule generation is a combination of kernel IF method and rough set theory. It will help to reduce the vagueness of patterns which are extracted using clustering techniques. Rough set theory has the ability to extract attribute sets without loss of the quality of approximation. Knowledge acquisition, decision support systems, and medical information are the fields using the concept of rough sets. The rule generation model can be used to improve the accuracy of data mining techniques. Initially, the clustering methods are applied on raw data and determine the indiscernibility relation. Dispensable and indispensable attributes will be generated from the groups. Finally, the discernibility matrix will be formed with all members. The research was experimented on IRIS and Habermann's survival dataset. The results were effective with less computation cost.

Leila Baccour et.al[9]., have presented some experiments for similarity and distance measures in IF sets. The concept of similarity measures in IF is used to compare patterns in systems such as handwritten Arabic sentences recognition, logic programming, and medical diagnosis. Similarity and distance measures are complement each other. If similarity measure increases then the distance measure decreases. The study has provided some explanations for definitions in similarity and distance measures. The similarity measures were experimented on squid

dataset which contains 1100 images of fishes. The concept of curvature scale space descriptors were used to identify edges of images. The eccentricity and the frequency of codes are the features used by the descriptors. The type – 1 fuzzy sets were applied to fuzzify descriptors and the eccentricity. The K Neural Network classifier was used to extract features from the dataset. The crisp feature vectors from the extracted features are used to form membership functions and fuzzification of features. The transformation of IF features will be performed using the fuzzy feature vectors. The study has experimented the similarity measures on Arabic sentence recognition. The results were shown the effectiveness of IF similarity measures.

Eulalia szmidt and Marta kookier[10] were discussed the classification of imbalanced and overlapping classes. The performance of classifier will be reduced by the imbalance and overlapping classes. Up – sampling and down – sampling are used to deal the imbalance problems. The authors were developed a method using fuzzy set approach to deal the problems. Confusion matrices are used to evaluate the behaviour of the classifier. The Atanassov's – IF (A- IF) is to describe the degree of freedom of fuzzy sets. The relation of A – Ifs can be used to convert the relative frequency distribution into an A – Ifs. The fuzzy sets and A – IF sets were applied on datasets. The performance of classifiers were also evaluated by the confusion matrices. Hesitation margins were also taken into account for the process of evaluation. Wine dataset and UCI ML repository were used for the experiments. The accuracy of A – IF a set has reached an average accuracy of 98% for wine dataset and fuzzy classifier has reached only 95%.

### 3. RESEARCH METHODOLOGY

The IFC is the extension of fuzzy clustering[23][24]. IFC algorithm has more flexible functionalities for large data corpus. C – means algorithm in IFC will produce optimal cluster centers and membership grades for each data. The membership grade will be useful to build inference system that represents each data members.

The figure 2 represents the processes that are involved in the research. Initially, data will be preprocessed and transformed into computable form. The pre – process activities are detailed in the next section. Data pre – process is one of the important factors for the generation of interesting patterns. The pre – processed data will be labeled to assess the feature selection process. The research has used manual labeling method for labeling the dataset. A converter program was

written in java to convert data into Extensible Markup Language (XML) format.

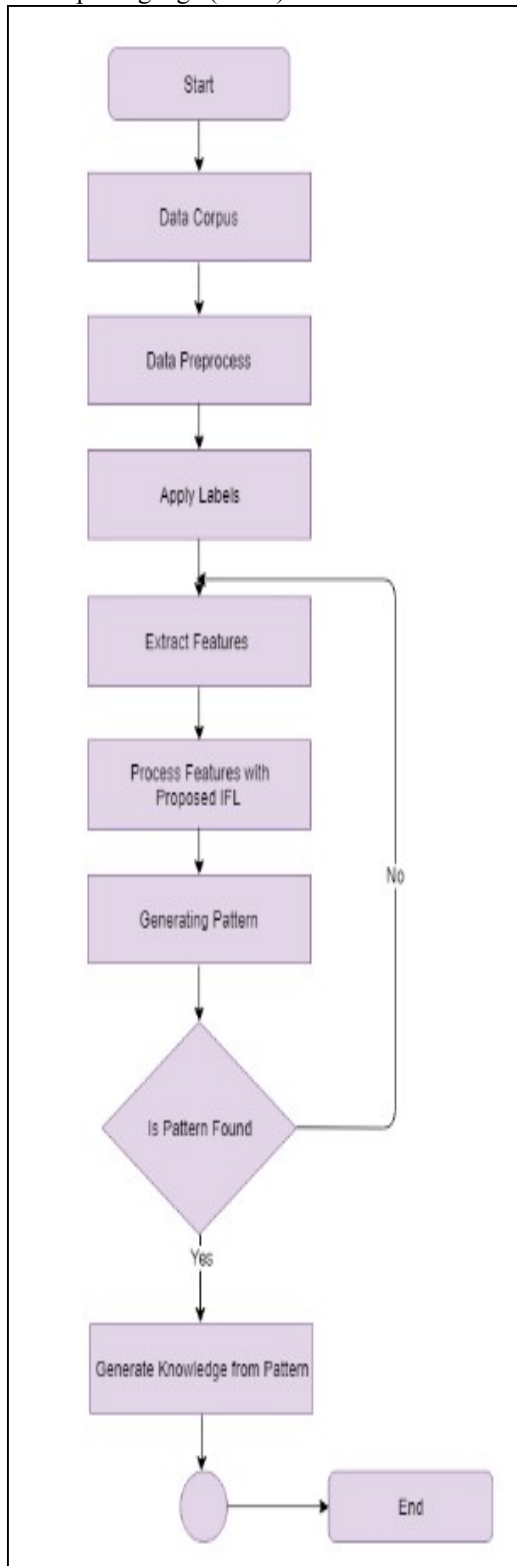


Figure 2 Flowchart – Pattern Extraction

XML is an easy to use format and reduce the complexities in pattern generation. XML tags are easily identifiable by any programming language. The research has employed Java Development Kit (JDK) 1.8 for the implementation of IFC C – means algorithm. The following figure 3 provides details about the procedure of extraction of patterns from normal dataset.

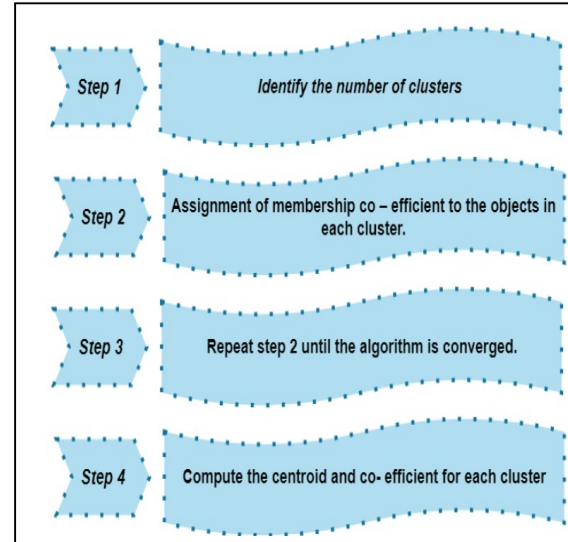


Figure 3 Procedures – Pattern Extraction for normal dataset

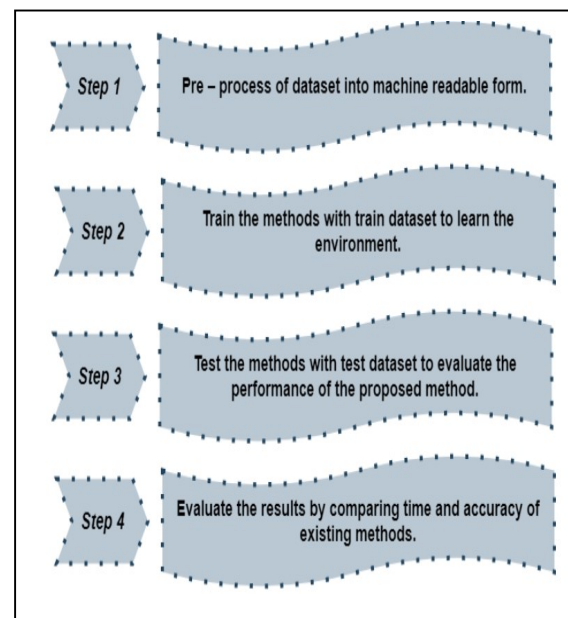


Figure 4 Procedures – Pattern Extraction for large dataset

Figure 4 represents the procedures for the extraction of pattern from large data corpus. The research has employed SVM, BIRCH, and Proposed IFC C – Means algorithm. The following snippet 1 - 7 show the IFC C – Means algorithm implementation in java.

*Snippet 1 – Initialize Cluster Membership*

```
int nCLIFC = U[0].length;
final Random Random(seedEnabled, randomSeed);
for (int i = 0; i < U.length; i++)
{
    float sum = 0;
    for (int j = 0; j < nCLIFC; j++)
    {
        U[i][j] = random.nextFloat();
        sum += U[i][j];
    }
    for (int j = 0; j < nClusters; j++)
    {
        U[i][j] /= sum;
    }
}
```

*Snippet 2 – Initialization of Random clusters*

```
while (clusterCreated < k)
{
    final int Candidate = random.next(nbPixels);

    if (!centerLocations.contains(clusterCandidate))
    {
        centerLocations.add(clusterCandidate);
        clusterCreated++;
    }
}
```

*Snippet 3 – Initialize Random Cluster Membership*

```
for (int i = 0; i < U.length; i++)
{
    float sum = 0;
    for (int j = 0; j < nCLIFC; j++)
    {
        U[i][j] = random();
        sum += U[i][j];
    }
    for (int j = 0; j < nCLIFC; j++)
    {
```

```
U[i][j] /= sum;
}
```

*Snippet 4 – Finding Closest Cluster*

```
for (int i = 0; i < k; i++)
{
```

```
    if (d < minimum_distance)
    {
        minimum_distance = d;
        closestCluster = i;
    }
}
```

*Snippet 5 – Finding Closest Cluster*

```
for (int i = 0; i < A.length; i++)
{
    D[i][B.length] = 0.0;
    for (int j = 0; j < B.length; j++)
    {
        double sum = 0;
        for (int k = 0; k < ns; k++)
        {
            final double d = D[i][k] - C[j][k];
            sum += d * d;
        }

        D[i][j] = sum;
        if (D[i][j] == 0.0)
        {
            D[i][B.length] = -1;
        }
        else
        {
            if (D[i][B.length] != -1)
            {
                D[i][B.length] += power(1 / D[i][j], exp);
            }
        }
    }
}
```

*Snippet 6 – Finding Exponential membership*

```
for (int i = 0; i < U.length; i++)
{
    for (int j = 0; j < nClusters; j++)
    {
        Um[i][j] = power(U[i][j], m);
        if (Float.isNaN(Um[i][j]))
        {
            Um[i][j] = 0;
        }
    }
}
```

## Snippet 7 – Defuzzyfy Clusters

```

int nCLIFC = U[0].length;
int [][] defU = new int [U.length][ nCLIFC];

for (int i = 0; i < U.length; i++)
{
    float max = -1;
    int max_Pos = -1;
    for(int j = 0; j < nCLIFC; j++)
    {
        if ( U[i][j] > max)
        {
            max = U[i][j];
            maxPos = j;
        }
    }

    defuzy[i][max_Pos] = 1;
}

return defuzy;
}

```

C – Means algorithm is similar to K –means but the performance of fuzzy C – Means is better than those algorithms. The reason for the better performance is the concept of fuzzy.

#### 4. DATA PREPARATION

Data preparation is an important phase in knowledge extraction. The pre – process activity will filter unwanted data and pass useful data to the classifier. Further, data will be transformed into computable form. Feature set will be selected from the transformed data and forward to the classifier. Figure 5 shows the detailed procedure of data preparation. Usually, large data set have more useless data, which forms complexity in the later part of research. Each process were carried out in different phases. Manual pre – process of data will take time but fruitful for generation of interesting patterns from dataset.

Data transformation were carried out for feature set extraction. The feature set is a collection of transformed data. The collection should be in a proper manner, so that the classifiers can understand and learn the environment. The feature sets were divided into two sets: training set and test set. Data in training and test sets are unique in nature. Special care was taken in segregating data into two divisions.

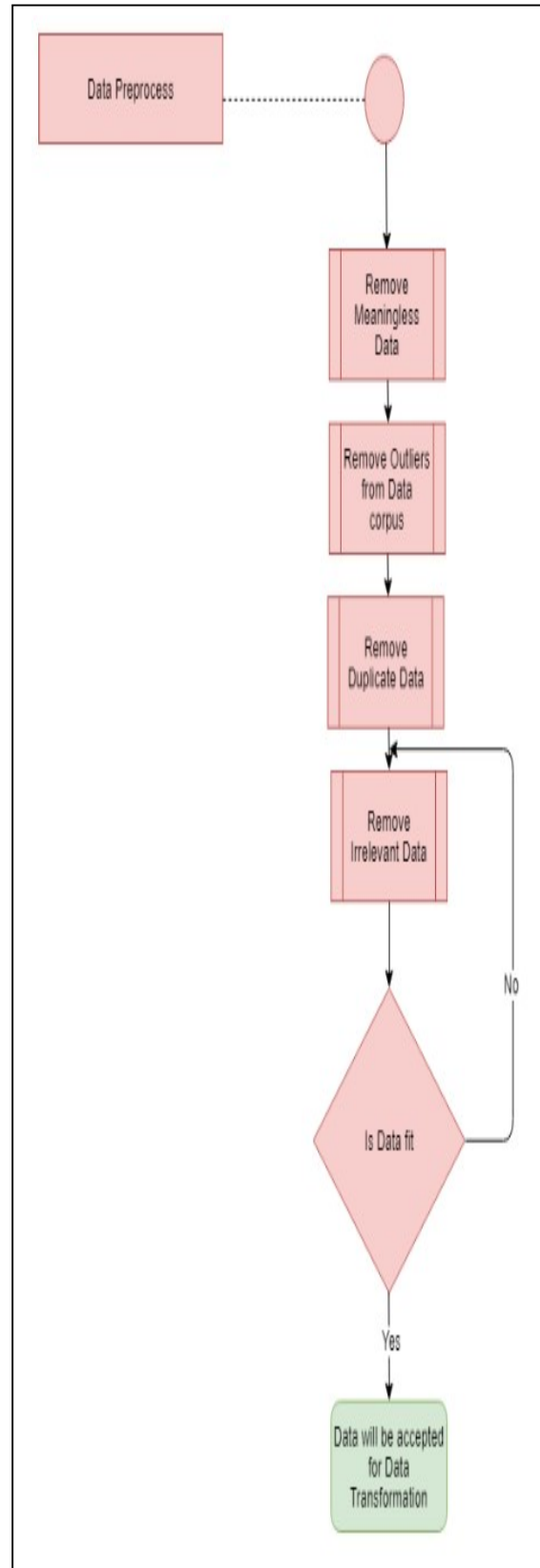


Figure 5 Procedures – Data Preparation



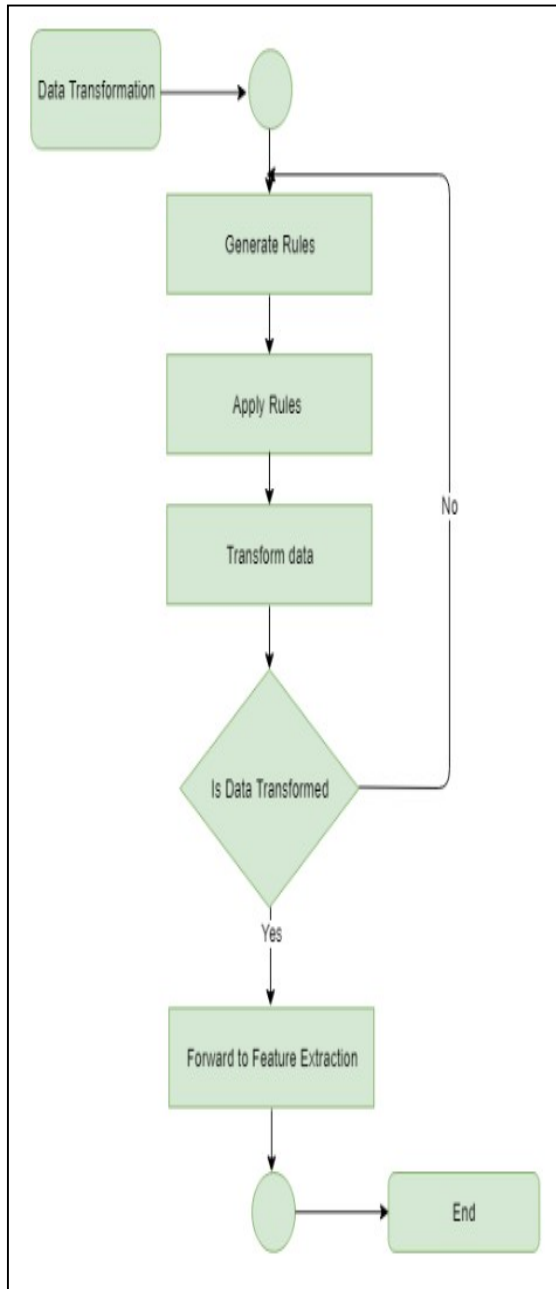


Figure 6 Data Transformation

Figure 6 shows the processes involved in the data transformation. The rules has to be generated to transform data into machine readable format. Data are complex in nature and difficult to read by ML methods. The transformation of data into computable format is the vital part of the research work. The research has carefully done the data preparation by following globally acceptable rules. After applying rules, data will be transformed into ML computable data and ready for extraction of features. Microsoft Excel 2007 and WEKA 3. 9.

3 were used for the preprocess of data. Figure 5 shows the initial phase of data preprocess. The pre – process activities were carried out using Microsoft Excel. Macros were developed to reject unwanted data, selection of useful data from data corpus. Figure 7 represents the extraction of feature sets from transformed data.

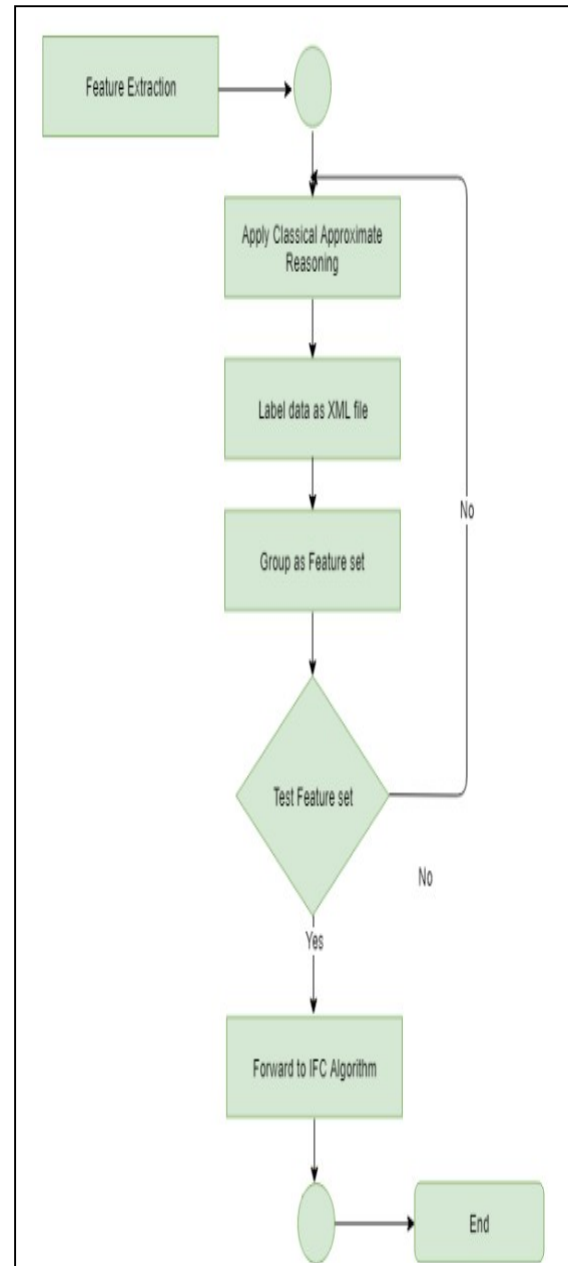


Figure 7 Extraction of Feature sets

The classical approximate reasoning was used for the extraction of feature sets. The concept of approximate reasoning will be used to select the appropriate features for the classifiers.

## 5. EXPERIMENT AND RESULTS

The methods were implemented in Java and Matlab in windows 10 operating systems. The WEBKB dataset found in the [www.cs.cmu.edu/~webkb/](http://www.cs.cmu.edu/~webkb/) used for the purpose of classification. The dataset were already classified useful for the evaluation of accuracy of the methods used in the research. BIRCH, SVM were the methods used in the research to compare proposed method. The objective of the research is to reduce the computation time and increase the prediction accuracy of the classification. The following table 1 shows the manual classification of WEBKB dataset.

Table 2 Classification of WEBKB

| Category   | Web pages |
|------------|-----------|
| Student    | 1641      |
| Faculty    | 1124      |
| Staff      | 137       |
| Department | 182       |
| Course     | 930       |
| Project    | 504       |
| Other      | 3764      |

Table 3 shows the training time of the methods. A maximum of 0.659 seconds consumed by SVM to classify staff data. A minimum of 0.231 seconds consumed by proposed method to produce project data. Figure 8 shows the relevant graph of training time of methods. All methods have taken more time as to learn the environment to produce effective results. The acronym “Stu” is students, “Fac” is Faculties, “Sta” is Staff, “Dep” is Department, “Cou” is Course, “Pro” is project, “Pro. Me” is Proposed Method.

Table 3 Training time (in Seconds) of WEBKB Dataset

| Category/<br>Algorithm | Stu  | Fac  | Sta  | Dep  | Cou  | Pro  |
|------------------------|------|------|------|------|------|------|
| <b>SVM</b>             | 0.78 | 0.65 | 0.68 | 0.54 | 0.62 | 0.58 |
| <b>BIRCH</b>           | 0.68 | 0.54 | 0.60 | 0.50 | 0.45 | 0.28 |
| <b>Pro.Me</b>          | 0.36 | 0.27 | 0.42 | 0.35 | 0.23 | 0.23 |

Approximate learning helps Intuistic fuzzy[26][27] to identify better features. Fuzzy c – means have options to produce results in less amount of time. The proposed method have less training time shows the performance of C – means. The approximate reasoning has helped the proposed method to learn the environment in short span of time with high accuracy.

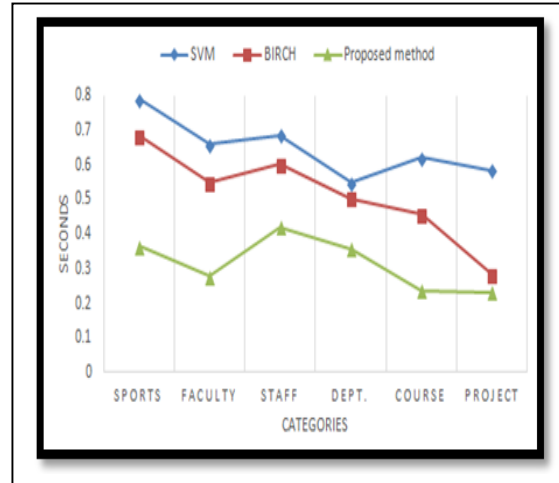


Figure 8 Training time (seconds) of WEBKB

Table 4 shows the training time of the methods. A maximum of 0.289 seconds consumed by BIRCH to classify staff data. A minimum of 0.101 seconds consumed by proposed method to produce course data. Figure 9 shows the relevant graph of testing time of methods. Proposed method took less time comparing to the existing methods show that the effective performance towards large dataset. Proposed method can classify data without using label as it learns the environment during training phase.

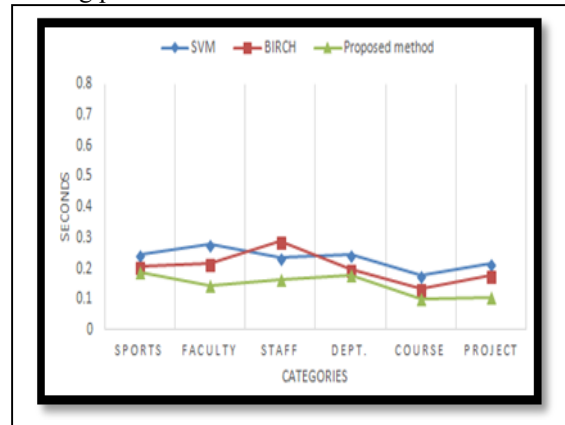


Figure 9 Testing time (seconds) of WEBKB

Table 4 Testing time (in Seconds) of WEBKB Dataset

| Category/<br>Algorithm | Stu  | Fac  | Sta  | Dep  | Cou  | Pro  |
|------------------------|------|------|------|------|------|------|
| <b>SVM</b>             | 0.24 | 0.27 | 0.23 | 0.24 | 0.17 | 0.21 |
| <b>BIRCH</b>           | 0.20 | 0.21 | 0.28 | 0.19 | 0.13 | 0.17 |
| <b>Pro.Me</b>          | 0.18 | 0.14 | 0.16 | 0.17 | 0.10 | 0.11 |



Table 5 shows the prediction accuracy of the methods during the training phase. The proposed method has scored minimum accuracy than existing methods. The proposed learner is a slow learner comparing to existing methods. The accuracy of SVM was better than other two methods. SVM will learn the environment in short span of time but consume more memory and computation time. The figure 10 shows the accuracy graph relevant to table 5.

Table 5 Prediction Accuracy of WEBKB Dataset  
(Training phase)

| Category/<br>Algorithm | Stu  | Fac  | Sta  | Dep  | Cou  | Pro  |
|------------------------|------|------|------|------|------|------|
| SVM                    | 87.8 | 81.3 | 82.5 | 76.5 | 81.8 | 83.9 |
| BIRCH                  | 81.6 | 78.4 | 84.2 | 79.3 | 75.2 | 84.3 |
| Pro.Me                 | 77.2 | 79.5 | 83.7 | 79.7 | 78.3 | 80.6 |

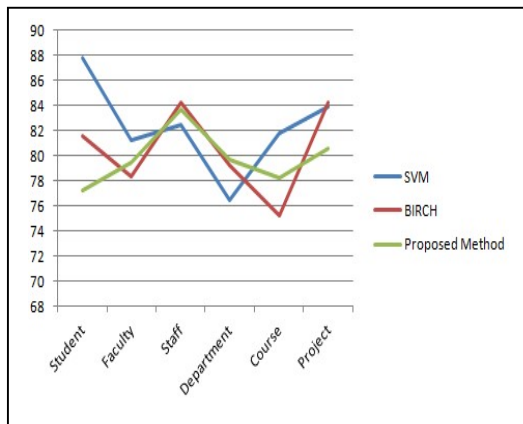


Figure 10 Prediction accuracy of methods (Training phase)

Table 6 shows the prediction accuracy of the methods used in the research. The proposed method outperformed the existing methods. It has scored more than 91 % for all the classification. SVM had below 90% and BIRCH got average of 90%. The figure 11 shows the accuracy graph relevant to table 6.

Table 6 Prediction Accuracy of WEBKB Dataset  
(Testing phase)

| Category/<br>Algorithm | Stu  | Fac  | Sta  | Dep  | Cou  | Pro  |
|------------------------|------|------|------|------|------|------|
| SVM                    | 89.3 | 88.1 | 87.4 | 79.5 | 84.6 | 87.6 |
| BIRCH                  | 91.3 | 90.6 | 89.3 | 90.5 | 91.3 | 88.9 |
| Pro.Me                 | 92.3 | 92.8 | 93.2 | 91.4 | 93.6 | 95.6 |

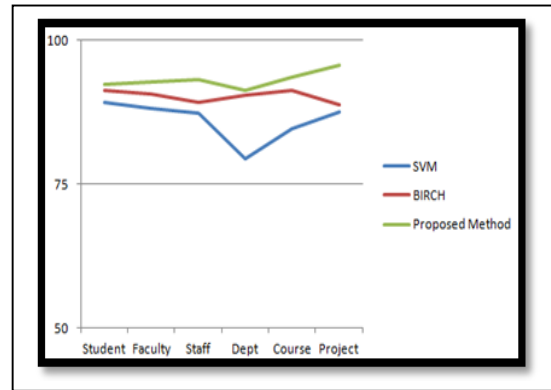


Figure 11 Prediction accuracy of methods (Testing phase)

Precision, Recall and F1 score are the evaluation metrics to know the retrieval capacity of the methods used for the information retrieval. We have evaluated the metrics for the methods used in the research. The following table 7 and 8 shows the precision and recall values of the methods. The proposed method have better recall and precision values comparing to existing methods. The F1 score combines both precision and recall values. Table 9 shows the F1 score of the methods, proposed method got better score than the existing methods. Figure 12 shows the F1 score relevant to the table 9

Table 7 Precision for WEBKB websites

| Category/<br>Algorithm | Stu | Fac | Sta | Dep | Cou | Pro |
|------------------------|-----|-----|-----|-----|-----|-----|
| SVM                    | 88  | 86  | 81  | 78  | 79  | 84  |
| BIRCH                  | 87  | 84  | 83  | 84  | 81  | 85  |
| Pro.Me                 | 91  | 89  | 88  | 86  | 83  | 80  |

Table 8 Recall for WEBKB websites

| Category/<br>Algorithm | Stu | Fac | Sta | Dep | Cou | Pro |
|------------------------|-----|-----|-----|-----|-----|-----|
| SVM                    | 87  | 86  | 79  | 94  | 84  | 86  |
| BIRCH                  | 86  | 82  | 84  | 85  | 87  | 87  |
| Pro.Me                 | 90  | 88  | 86  | 84  | 90  | 86  |

Table 9 F1 score of WEBKB

| Metrics | Precision | Recall | F1 Score |
|---------|-----------|--------|----------|
| SVM     | 82.6      | 86     | 84.2     |
| BIRCH   | 84        | 85.1   | 84.5     |
| Pro.Me  | 86.1      | 87.3   | 86.7     |

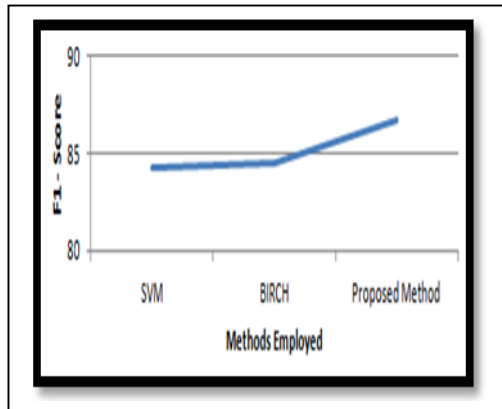


Figure 12 F1 score of methods

### Conclusion

The proposed method combines approximate learning and intuitionistic fuzzy logic to classify large data corpus. Approximate learning is one of the pioneer techniques to predict accurate results from complex and uncertain data. Intuitionistic fuzzy has the ability to produce effective results in short duration of time. The combination of these methods has improved the percentage of accuracy and time to classify large dataset. The WEBKB is the large data corpus contains more than 8000 data of 4 different Universities. The results and experiment has shown the effective performance of the proposed method. The proposed method have scored 86% of F1 score shows the better retrieval capacity and took less computation time and achieved more than 90% of prediction accuracy.

### REFERENCES

- [1] Changyu Liu , Wenyan Gan , Tao Wu, A comparative study of cloud model and extended fuzzy sets, Proceedings of the 5th international conference on Rough set and knowledge technology, October 15-17, 2010, Beijing, China
- [2] Feng Feng , Yongming Li , Violeta Leoreanu-Fotea, Application of level soft sets in decision making based on interval-valued fuzzy soft sets, Computers & Mathematics with Applications, v.60 n.6, p.1756-1767, September, 2010
- [3] Dongrui Wu , Jerry M. Mendel, Perceptual reasoning for perceptual computing: a similarity-based approach, IEEE Transactions on Fuzzy Systems, v.17 n.6, p.1397-1411, December 2009
- [4] Marijana Gorjanac Ranitović , Aleksandar Petojević, Short communication: Lattice representations of interval-valued fuzzy sets, Fuzzy Sets and Systems, 236, p.50-57, February, 2014
- [5] Jin Han Park , Jong Seo Park , Young Chel Kwun, On fuzzy inclusion in the interval-valued sense, Proceedings of the Second international conference on Fuzzy Systems and Knowledge Discovery, August 27-29, 2005, Changsha, China.
- [6] Peerasak Intarapaiboon, "Text classification using similarity measures on intuitionistic fuzzy sets", ScienceAsia (2016), volume 42, PP. 52 - 60.
- [7] S.V.Aruna Kumar and B.S.Harish, " A modified intuitionistic fuzzy clustering algorithm for medical image segmentation", Journal of Intelligence System, 2017, PP.1 - 15.
- [8] Kuo - Ping Lin, Kuo - Chen Hung, and Ching - Lin Lin, " Rule generation based on novel kernel intuitionistic fuzzy rough set model," IEEE Access, Volume 6, 2018.
- [9] Leila Baccour, Adel M. Alimi and Robert I.John, " Intuitionistic fuzzy similarity measures and their role in classification", Journal of Intelligence systems, 2015. PP. 1 - 17
- [10] E. Szmidt and M. Kukier: Atanassov's Intuitionistic Fuzzy Sets in Classification of Imbalanced and Overlapping Classes, Studies in Computational Intelligence (SCI) 109, 455-471 (2008).
- [11] R. Roy , P. K. Maji, A fuzzy soft set theoretic approach to decision making problems, Journal of Computational and Applied Mathematics, v.203 n.2, p.412-418, June, 2007
- [12] Feng Feng , Young Bae Jun , Xiaoyan Liu , Lifeng Li, An adjustable approach to fuzzy soft set based decision making, Journal of Computational and Applied Mathematics, v.234 n.1, p.10-20, May, 2010
- [13] Danqing Xu , Yanan Fu , Junjun Mao, Dynamic Analysis of IVFSs Based on Granularity Computing, Proceedings of the 14th International Conference on Rough Sets, Fuzzy Sets, Data Mining, and Granular Computing, October 11-14, 2013
- [14] Xiuqin Ma , Norrozila Sulaiman , Hongwu Qin , Tutut Herawan , Jasni Mohamad Zain,

- A new efficient normal parameter reduction algorithm of soft sets, *Computers & Mathematics with Applications*, v.62 n.2, p.588-598, July, 2011
- [15] Guiwu Wei , Hongjun Wang , Xiaofei Zhao , Rui Lin, Hesitant triangular fuzzy information aggregation in multiple attribute decision making, *Journal of Intelligent & Fuzzy Systems: Applications in Engineering and Technology*, v.26 n.3, p.1201-1209, May 2014
- [16] Degang Chen , E. C. C. Tsang , Daniel S. Yeung , Xizhao Wang, The parameterization reduction of soft sets and its applications, *Computers & Mathematics with Applications*, v.49 n.5-6, p.757-763, April, 2005
- [17] Xiuqin Ma , Norrozila Sulaiman, An interval-valued fuzzy soft set approach for normal parameter reduction, *Proceedings of the 13th international conference on Rough sets, fuzzy sets, data mining and granular computing*, June 25-27, 2011, Moscow, Russia
- [18] Dongrui Wu , Jerry M. Mendel, A comparative study of ranking methods, similarity measures and uncertainty measures for interval type-2 fuzzy sets, *Information Sciences: an International Journal*, v.179 n.8, p.1169-1192, March, 2009
- [19] Sabir Hussain , Bashir Ahmad, Some properties of soft topological spaces, *Computers & Mathematics with Applications*, v.62 n.11, p.4058-4067, December, 2011
- [20] Jinyan Wang , Minghao Yin , Wenxiang Gu, Soft polygroups, *Computers & Mathematics with Applications*, v.62 n.9, p.3529-3537, November, 2011
- [21] Aslıhan Sezgin , Akın Osman Atagün, Soft groups and normalistic soft groups, *Computers & Mathematics with Applications*, v.62 n.2, p.685-698, July, 2011
- [22] Barbara Pekala, Operations on Interval Matrices, *Proceedings of the international conference on Rough Sets and Intelligent Systems Paradigms*, June 28-30, 2007, Warsaw, Poland
- [23] Peide Liu, A weighted aggregation operators multi-attribute group decision-making method based on interval-valued trapezoidal fuzzy numbers, *Expert Systems with Applications: An International Journal*, v.38 n.1, p.1053-1060, January, 2011
- [24] Ranjit Biswas, ON i-v FUZZY SUBGROUPS, *Fundamenta Informaticae*, v.26 n.1, p.1-9, January 1996
- [25] Feng Feng , Young Bae Jun , Xianzhong Zhao, Soft semirings, *Computers & Mathematics with Applications*, v.56 n.10, p.2621-2628, November, 2008
- [26] Jun Ye, Multicriteria fuzzy decision-making method based on a novel accuracy function under interval-valued intuitionistic fuzzy environment, *Expert Systems with Applications: An International Journal*, v.36 n.3, p.6899-6902, April, 2009
- [27] Adam Niewiadomski , Michał Bartyzel, Elements of the type-2 semantics in summarizing databases, *Proceedings of the 8th international conference on Artificial Intelligence and Soft Computing*, June 25-29, 2006, Zakopane, Poland.