

INDIRA GANDHI NATIONAL OPEN UNIVERSITY

Master of Computer Applications (MCA)

MCSL-223: Computer Networks and Data Mining Lab*Section 2: Data Mining*

Semester: II **Sessions:** 10 **Exercises:** 23 **Source:** IGNOU lab manual, list of lab exercises

Instructions

1. Use the open source WEKA toolkit (version 3.8 or later). The sample datasets are in the data folder of the installation.
2. For every run, record the exact parameter values, the run information block and the numbers you can verify by hand (support, confidence, entropy, accuracy, kappa, sum of squared errors).
3. Interpret every output; a screenshot alone does not answer an exercise.

Session 1: WEKA, datasets and ARFF files

Q1. Download and install WEKA. Navigate the various options available in WEKA. Explore the available datasets in WEKA. Load various datasets and observe the following:

1. List the attribute names and their types
2. No. of records in each dataset
3. Identify the class attribute (if any)
4. Plot Histogram
5. Determine the no. of records for each class.
6. Visualize the data in different dimensions.

Q2. Create your own EXCEL file. Convert the EXCEL file to .csv format and prepare it as .arff file.

Q3. Try to create your own datasets.

Q4. Preprocess and classify Customer, Agriculture, Weather, Whole-sale Customers or the datasets of your own choice from the [UCI Machine Learning Repository](#).

Session 2: Basic preprocessing

Q5. Perform the basic pre-processing operations on data relation such as removing an attribute and filter attribute bank data.

Q6. Demonstrate the preprocessing mechanism on the following datasets:

1. student.arff
2. labor.arff
3. contactlenses.arff

Session 3: Unsupervised filters and Apriori

Q7. Perform the following:

- Explore various options available in WEKA for preprocessing data and apply unsupervised filters like Discretization, Resample-filter etc. on various datasets.
- Load weather, nominal, Iris, Glass datasets into WEKA and run Apriori algorithms with different support and confidence values.
- Study the rules generated.
- Apply different discretization filters on numerical attributes and run the Apriori association rule algorithm. Study the rules generated.
- Derive interesting insights and observe the effect of discretization in the rule generation process.

Q8. Implement the Apriori Algorithm to find the association rules in contactlenses.arff dataset.

Session 4: FP-Growth and Apriori parameter studies

Q9. Find the frequent patterns using FP-Growth algorithm on contactlenses.arff and test.arff datasets.

Q10. Generate association rules using Apriori algorithm with Bank.arff relation

1. Set minimum support range as 20% to 100%, incremental decrease factor as 5% and confidence factor as 80% and generate 5 rules.
2. Set minimum support as 10%, delta 5%, minimum lift as 150% and generate 4 rules.

Q11. Generate association rule for the credit card promotion dataset using Apriori algorithm with the support range 40% to 100%, confidence as 10%, incremental decrease as 5% and generate 6 rules.

Session 5: Reading Apriori output

Q12. Perform the following:

- Use contactlenses.arff and load it into WEKA. Check that all attributes are nominal (categorical).
- Change to the Associate Panel. Select Apriori as associator. After pressing the start button, Apriori starts to build its model and writes its output into the output field. The first part of the output (Run information) describes the options that have been set and the data set used. Make sure you understand all the data reported.

- The rules that have been generated are listed at the end of the output. By default, only the 10 most valuable rules according to their confidence level are shown. Each rule consists of some attribute values on a left hand side of the arrow, the arrow sign and the right hand side list of attribute values. Right of the arrow sign are the predicted attribute values. Rules have certain support and confidence values. The number before the arrow sign is the number of instances the rule applies to. The number after the arrow sign is the number of instances predicted correctly. The number in brackets after conf: is the confidence of the rule. Analyse the rules mined from the data set. What are their confidence and support values? Examine the number of large itemsets and make sure you understand how this data has been calculated (check that the values you would get manually are correct).

Session 6: Apriori on the zoo dataset and regression

Q13. Perform the following:

- Use zoo.arff dataset and load it into WEKA. Examine the attributes and make sure you understand their meaning. Are all attributes nominal?
- In the preprocess area, deselect the animal and legs attributes. The animal attribute is the name of the animal, and is not useful for mining. The legs attribute is numeric and cannot be used directly with Apriori. Alternatively, you can try to use the Discretize Filter to discretize the legs attribute.
- After deselecting the attributes, use the Apply Filters button to generate a working relation that removes those attributes. Notice how the working relation changes, and has fewer attributes than the base relation.
- First, try using the Apriori algorithm with the default parameters. Record the generated rules.
- Vary the number of rules generated (click on the command that you are running). Try 20, 30, and so on. Record how many rules you have to generate before generating a rule containing type=mammal.
- Vary the maximum support until a rule containing type=mammal is the top rule generated. Record the maximum support needed.
- Select one generated rule that was interesting to you. Why was it interesting? What does it mean? Check its confidence and support: are they high enough?
- Suggest one improvement to the Apriori implementation in WEKA that would have made this data mining lab easier to accomplish.

Q14. Demonstrate to predict the Numerical Values in the given Data Set is using Regression Methods.

Session 7: Classification: logistic regression,

decision tree, Naive Bayes, k-NN, SVM

Q15. Demonstrate the classification rule process on the student.arff, employee.arff and labor.arff datasets using the following algorithms:

1. Logistic Regression
2. Decision Tree
3. Naive Bayes

Q16. Demonstrate the classification rule process on the student.arff, employee.arff and labor.arff datasets using the following algorithms:

1. K-Nearest Neighbour
2. SVM

Session 8: Decision trees, rules, entropy, kappa and ROC

Q17. Perform the following:

- Demonstrate performing classification on various data sets.
- Load each dataset into WEKA and run ID3, J48 classification algorithm.
- Study the classifier output. Compute entropy values, Kappa statistic.
- Extract if-then rules from the decision tree generated by the classifier.
- Observe the confusion matrix.

Q18. Perform the following:

- Load each dataset into WEKA and perform Naive Bayes classification and k-Nearest Neighbour classification.
- Interpret the results obtained.
- Plot ROC Curves.
- Compare classification results of ID3, J48, Naive Bayes and k-NN classifiers for each dataset, and deduce which classifier is performing best and poor for each dataset and justify.

Session 9: k-means clustering

Q19. Demonstrate Clustering features in Large Databases with noise.

Q20. Implement simple K-Means Algorithm to demonstrate the clustering rule on the following datasets:

1. iris.arff
2. student.arff

Q21. Perform the following:

- Load each dataset into WEKA and run simple k-means clustering algorithm with different values of k (number of desired clusters).
- Study the clusters formed.
- Observe the sum of squared errors and centroids, and derive insights.
- Explore other clustering techniques available in WEKA.
- Explore visualization features of WEKA to visualize the clusters.
- Derive interesting insights and explain.

Session 10: Hierarchical and density-based clustering

Q22. Implement Hierarchical Clustering Algorithm to demonstrate the clustering rule process in the following datasets:

1. `employee.arff`
2. `student.arff`

Q23. Implement Density based Clustering Algorithm to demonstrate the clustering rule process on dataset `employee.arff`.

End of question paper. Worked solutions for every exercise: syntax.theether.in/mcs1-223/section-2/