

Artificial System for Prediction of Student's Academic Success from Tertiary Level in Bangladesh

Dr. Linkon Chowdhury¹ and Shahana Yeasmin²

¹ BGC Trust University Bangladesh

Received: 16 December 2011 Accepted: 31 December 2011 Published: 15 January 2012

Abstract

Every year a large scale of students in Bangladesh enrol in different Universities in order to pursue higher studies. With the aim to build up a prosperous career these students begin their academic phase at the University with great expectation and enthusiasm. However among all these enthusiastic and hopeful bright students many seem to become successful in their academic career and found to pursue the higher education beyond the undergraduate level. The main purpose of this research is to develop a dynamic academic success prediction model for universities, institutes and colleges. In this work, we first apply chi square test to separate factors such as gender, financial condition and dropping year to classify the successful from unsuccessful students. The main purpose of applying it is feature selection to data. Degree of freedom is used to P-value (Probability value) for best predictors of dependent variable. Then we have classify the data using the latest data mining technique Support Vector Machines(SVM).SVM helped the data set to be properly design and manipulated. After being processed data, we used the MATH LAB for depiction of resultant data into figure. After being separation of factors we have had examined by using data mining techniques Classification and Regression Tree (CART) and Bayes theorem using knowledge base. Proposition logic is used for designing knowledge base. Bayes theorem will perform the prediction by collecting the information from knowledge Base. Here we have considered most important factors to classify the successful students over unsuccessful students are gender, financial condition and dropping year. We also consider the sociodemographic variables such as age, gender, ethnicity, education, work status, and disability and study environment that may influence persistence or academic success of students at university level. We have collected real data from Chittagong University Bangladesh from numerous students. Finally, by mining the

Index terms— Chi Value, P-value, CART, SVM, Bayes theorem, Degree of freedom.

1 Introduction

Student retention is an indicator of academic performance and enrolment management of the university. Increasing student retention or persistence is a long term goal in all academic institution. High rates of student attrition have been a

Author ? : Chittagong, Bangladesh. E-mail : linkon_cse_cu@yahoo.com Author ? : Chittagong, Bangladesh. E-mail : bubbly.csecu@gmail.com concern for the past several years at the Chittagong University Institute of Computer Science and Engineering. Level of retention rate goes higher by various reasons and one the most important reasons are government funding, scholarship in the tertiary education environment. There are both

academic and administrative staff are under pressure to come up with strategies that could increase retention rates on their courses and programs. The lowest student retention rates at all institutions of higher education are first-year students, who are at greatest risk of dropping out in the first term or semester of study or not completing their programs or degree. Therefore most retention studies address the retention of first-year students. Consequently, the early identification of vulnerable students who are prone to drop their courses is crucial for the success of any retention strategy. This would allow educational institutions to undertake timely and pro-active measures. Once identified, these at risk students can be then targeted with academic and administrative support to increase their chance of staying on the course.

Data from 2004 to 2010, covering over 200 enrolled students stored in the Computer Science and Engineering was used to perform a quantitative analysis of study outcome. There were three main types of objectives conducted in this study. The first approach is descriptive which is concerned with the nature of the dataset such as the frequency table and the relationship between the attributes obtained using cross tabulation analysis (contingency tables). In addition, feature selection is conducted to determine the importance of the prediction variables for modelling study outcome. The third type of data mining approach, i.e. predictive data mining is conducted by using two different types of classification trees. The classification tree models have some advantages. Secondly, the classification tree models are non-parametric and can capture nonlinear relationships and complex interactions between predictors and dependent variable. We decided not to use other data mining techniques such as neural networks and support vector machines even though in some cases they could achieve higher accuracy, because their structure is not transparent and usually described as a black box. It is also difficult to explain their results and how they work to a user who would like S value) for best predictors of dependent variable. Then we have classify the data using the latest data mining technique Support Vector Machines(SVM).SVM helped the data set to be properly design and manipulated. After being processed data, we used the MATH LAB for depiction of resultant data into figure. After being separation of factors we have had examined by using data mining techniques Classification and Regression Tree (CART) and Bayes theorem using knowledge base. Proposition logic is used for designing knowledge base. Bayes theorem will perform the prediction by collecting the information from knowledge Base. Here we have considered most important factors to classify the successful students over unsuccessful students are gender, financial condition and dropping year. We also consider the sociodemographic variables such as age, gender, ethnicity, education, work status, and disability and study environment that may influence persistence or academic success of students at university level. We have collected real data from Chittagong University Bangladesh from numerous students. Finally, by mining the data, the most important factors for student success and a profile of the typical successful and unsuccessful students are identified.

2 Global

3 C

to apply them to a new set of data. Finally, a comparison between these models was conducted to determine the best model for the dataset. The population of this study was entering CSE in the firstyear from 2004 through 2009. These were full-time degree-seeking students. Entering in the first-year session were excluded from the study for a few reasons. The independent or predictor variables fell into three main categories. These were the students' personal information, high school background. The personal information recorded for each student was: ? Ethnicity(Bengali, Marma, Chakma and Others)? Gender ? Age ? Financial Status ? Work Status ? Disability II.

4 Collected statistical data

As part of the data-understanding phase we carried out the data on the table 1 and table 2. The Table 2 reports the results. Based on the results shown majority of Information Systems students are female (over 38%). However, percentage of female students who successfully complete the course are higher (41%) which suggests that female students are more likely to [3] pass the course than their male counterpart. When it comes to age over 26% of students are above 24. This age group is also more likely to fail the course because their percentage of students who failed the course in this age group (11.9%) is higher than their overall participation in the student population (26.2%). Statistical data on 42 students: Students with it are more likely to fail than those without it. There are huge differences in percentage of students who successfully completed [4] the course depending on their ethnic origin. A substantial number of students (over 55%) have financial support more vulnerable than the other two categories in this variable.

5 III.

6 Support vector machines

Support Vector Machine (SVM) is one of the latest clustering techniques which enables machine learning concepts to amplify predictive accuracy in the case of axiomatically diverting data those are not fit properly. It uses inference space of linear functions in a high amplitude feature space, trained with a learning algorithm. It works

by finding a hyper plane that linearly separates the training points, in a way such that each resulting subspace contains only points which are very similar. First and foremost idea behind Support Vector Machines (SVMs) is that it constituted by set of similar supervised learning. An unknown tuple is labeled with the group of the points that fall in the same subspace as the tuple. Earlier SVM was used for Natural Image processing System (NIPS) but now it becomes very popular is an active part of the machine learning research around the world. It is also being used for pattern classification and regression based applications. The foundations of Support Vector Machines (SVM) have been developed by V.Vapnik.

SVM is very effective in various data and information classification process. An expert should bear in mind two important factors for implementing SVM, these two factors or techniques are mathematical programming and kernel functions. Kernel methods leads or portrayal data into colossal amplitude margins in the anticipation that in this colossal amplitude margin the data could become more easily separated or better structured. Mathematical Programming refers the conception of the Linear programming for the best fit of Hyper plane. The word programming means to plans or make a time table for regular work. Integer Linear programming (ILP) which is the part of linear programming is very useful analytical and engineering tools to get an optimal solution .The parameters are found by solving a quadratic programming problem with linear equality and inequality constraints; rather than by solving a nonconvex, unconstrained optimization problem. The flexibility of kernel functions allows the SVM to search a wide variety of hypothesis spaces. The for-most reasons of using SVM are to select the proper Support Vectors for the data classification. The figure 1 shows a graphical view of Support Vectors selection of the process. All hypothesis space help to identify the Maximum Margin Hyper plane (MMH) which enables to classify the best and almost correct data the following figure shows the process of SVMs selection from large amount of SVMs. We can calculate the weight boundary maximum margin by using the following equation:

Another interesting question is why maximum margin? There are some good explanations which include better empirical performance. Another reason is that even if we've made a small error in the location of the boundary this gives us least chance of causing a misclassification. The other advantage would be avoiding local minima and better classification. The goals of SVM are separating the data with hyper plane and extend this to non-linear boundaries using kernel trick [8] [11]. For calculating the SVM we see that the goal is to correctly classify all the data. For mathematical calculations we have, In this equation x is a vector point and w is weight and is also a vector. So to separate the data a should always be greater than zero. Among all possible hyper planes, SVM selects the one where the distance of hyper plane is as large as possible. If the training data is good and every test vector is located in radius r from training vector. Now if the chosen hyper plane is located at the farthest possible from the data [12]. This desired hyper plane which maximizes the margin also bisects the lines between closest points on convex hull of the two datasets. Thus we have a , b & c .

7 Classification and regression tree (cart)

Classification and Regression Tree (CART) has many advantages over classification methods. It is naturally non-parametric. CART can handle [5][6] numerical data that are highly skewed. It eliminates analyst time, which would otherwise be spent determining whether variables are normally distributed. CART identifies "splitting" variables based on a complete search of all possibilities. CART is able to search all possible variable splitters, even in problems with many hundreds of possible predictors as well as dealing missing variable. Finally, CART trees are moderately simple and non-statisticians.

The purpose of an analysis based on classification tree is to find out the factors that contribute the separation of successful and unsuccessful students among the all recorded data. When the classification tree is generated it is possible to calculate the probability of each student's outcome. If once the classification tree is formed, it is possible to predict the study outcome for newly enrolled student. In each tree node the percentages for successful and unsuccessful students is given and also absolute size of the node. The variable names above the node are the predictor's criteria that make the split for the node according to classification and regression tree. Each node is split according to predictor criteria. The searching is stops when the split with the largest improvement in goodness of fit. Possible predictor variables in the dataset were included in the classification tree in splits process, which is detected in feature selection criteria. We used two stopping criteria in the training process: 1. A maximum tree depth has been proceeding into 3 levels for CHAID tree. 2. A maximum tree depth has been proceeding into 4 levels for CART tree.

Lastly for each classification tree we have assigned different costs to the classification outcomes i.e.: classification matrix. It is one of the processes of increasing the correctly classified student's outcomes. We categories the student's outcome according to Pass and Fail in the academic courses

In this research we explore the concepts and technique of SVM to classify the data collected in our experiments. We have approximately collected twelve hundreds (1200) data from University of Chittagong where about thirty (30000) thousands students are studying. To assess these large amounts of data we have found that SVM is very efficient and exact technique in our proceedings. By imposing the SVM, we have mapped the data to meaning full forty two (42) data which are shown in the figure 3 and 4. In figure ?? we have depicts that the reasons for the age where mainly focused on the pivotal age of greater or less than twenty four. Fig. ?? : Data classification for age group using SVM From the figure above we can easily measure the Maximum Margin Hyper plane (MMH). At MMH the resultant outcome of age group using SVM is determined. We have also used the MATHLAB to

accelerate the accuracy of the implementation. In the same process we have had accomplished our design and implementation for the financial support data using the methodology of SVM. Step1: First insert the observed value in each cell of observable table. Inserted value collected from record.

Step 2: Calculate expected value for every cell of the describing table ?? Step 3: calculating chi value for every cell using the following formula: $\chi^2 = (\text{observed value} - \text{expected value})^2 / \text{expected value}$

Step 4: calculate total chi -value for domain using the following formula $(\text{observed value} - \text{expected value})^2 / \text{expected value}$

Step 5: calculating degree of freedom using following rule Degree of freedom, $df = (\text{No.of.rows} - 1) * (\text{No.of.columns} - 1)$

Step 6: calculate p-value (probability value) using following method in Ms Excel $P\text{-value} = \text{CHIDIS}(\text{Chi value}, df)$ Step 1: consider the domain is financial support in which category1=Yes category2=No, Option1=Pass and Option2=Fail. The value of every cell collects from database.

Step 2: calculating expected value for each cell using describing formula

Step 3: calculating chi value for every cell using the describing formula:

Step 4: calculate total chi -value for domain $\chi^2 = 0.283 + 0.631 + 0.343 + 0.764 = 2.021$ Overall chi-value for every domain in Pass and Fail Category: Table ?? : Individual Chi-value for each category

Step 5: calculating degree of freedom using following the rule Degree of freedom $df = (2-1) * (2-1) = 1$

Step 6: calculate p-value (probability value) using in Ms Excel $P\text{-value} = \text{CHIDIS}(2.021, 1) = 0.155$

Factors selection is an important process to assess the prediction of dropout of the students. The prediction has relate the variable that determines the rate of success The number of predictor variables is not so large and we don't have to select the subset of variables for further analysis which is the main purpose of applying feature selection to data. However, feature selection could be also used as a pre-processor for predictive data mining to rank predictors according to the strength of their relationship with dependent or outcome variable. During the factors selection process no specific form of relationship, neither linear nor nonlinear, is assumed. The outcome of the factors selection would be a rank list of predictors according to their importance for further analysis of the dependent variable with the other methods for regression and classification. Here the figure below shows the relative outcome of the predictor's value. In all three cases, i.e. for all three definitions of the dependent variable, if the top 4 variables are selected, we get the same list of predictors. Therefore we can conclude that the list of important predictors is quite robust to changes in the study outcome definition. We may proceed into the next step using the top 4 variables: 1. Financial Support 2. Age Group 3. Gender 4. Disabilities From Table 4, P-values from the last column only the first three chi-square values are significant at 10% level. Though the results of the feature selection suggested continuing analysis with only the subset of predictors, which includes Financial Support, Age group and Gender, we have included all available predictors in our classification tree analysis. We follow an advice given in Luan & Zhao (2006) who suggested that even though some variables may have little significance to the overall prediction outcome, they can be essential to a specific record. In our work we have described a simple method of probabilistic inference that is, the computation from observed evidence of posterior probabilities for query propositions. We have used the joint probability as the knowledge base from which answer to all question may be derived. We have had built the knowledge base by considering two Boolean variables. The table 7 is an example of two valued propositional logic which is the bases of knowledge base representation: Table ??: Concepts of propositional logic to design a Knowledge Base using the proposition of Boolean events A, B and C.

Based on table 7, we have designed the knowledge base (Joint probability distribution) for our research activity. Here we have considered those events which have true (one or 1) Boolean values. Table ?? is an example of knowledge base for events A, B and C: Bayes' theorem and conditional probability are opposite to each other. Given two dependent events A and B. The conditional probability of P (A and B) or P (B/A) will be P (A and B)/P (A). Related to this formula a rule is developed by the English Presbyterian minister Thomas Bayes (1702-61).According to the Bayes rule it is possible to determine the various probabilities of the first event given the outcome of the second event in a sequence of two events.

The conditional probability:

(1)

The equation (1) will help to find out the probabilities of B after being occurrences of the A. we get the Bayes' theorem for these two events as follows:

(2) If there are more events like A1, A2, and B1, B2.In this case the Bayes theorem to determine the probability of A1 based on B1will be as follows:

Consideration Classification and Regression Tree (CART):

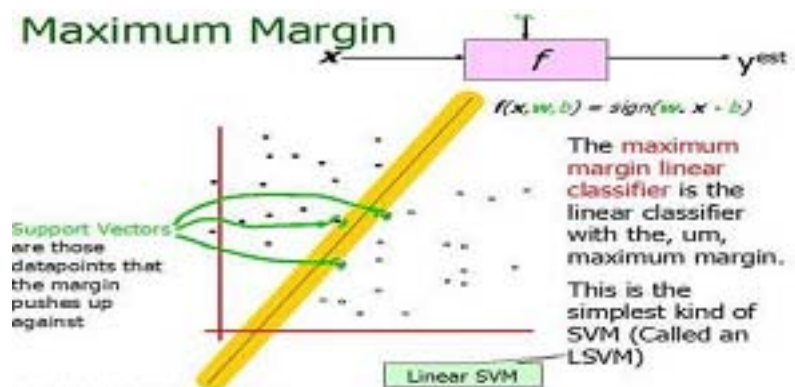
For Node 4: IF Financial Support="Yes" AND Age $\leq$ 24 THEN "PASS"= $(18/23) * (9/15) = 0.49$

Now applying the Bayes theorem on table 5 we have got the following outcomes:

If one student pass based on his financial condition is "Yes" and age $\leq$ 24 then $P(\text{Pass} | \text{Financial condition} = \text{"Yes"} \wedge \text{age} \leq 24) = P(\text{Pass} \wedge \text{Financial condition} = \text{"Yes"} \wedge \text{age} \leq 24) / P(\text{Financial condition} = \text{"Yes"} \wedge \text{age} \leq 24)$ $P(\text{Pass} \wedge \text{Financial condition} = \text{"Yes"} \wedge \text{age} \leq 24) = 0.279$ $P(\text{Financial condition} = \text{"Yes"} \wedge \text{age} \leq 24) = 0.125 + 0.279 = 0.404$ $P(\text{Pass} | \text{Financial condition} = \text{"Yes"} \wedge \text{age} \leq 24) = 0.125 / 0.404 = 0.31$



Figure 1:



1

Figure 2: GlobalFigure 1 :

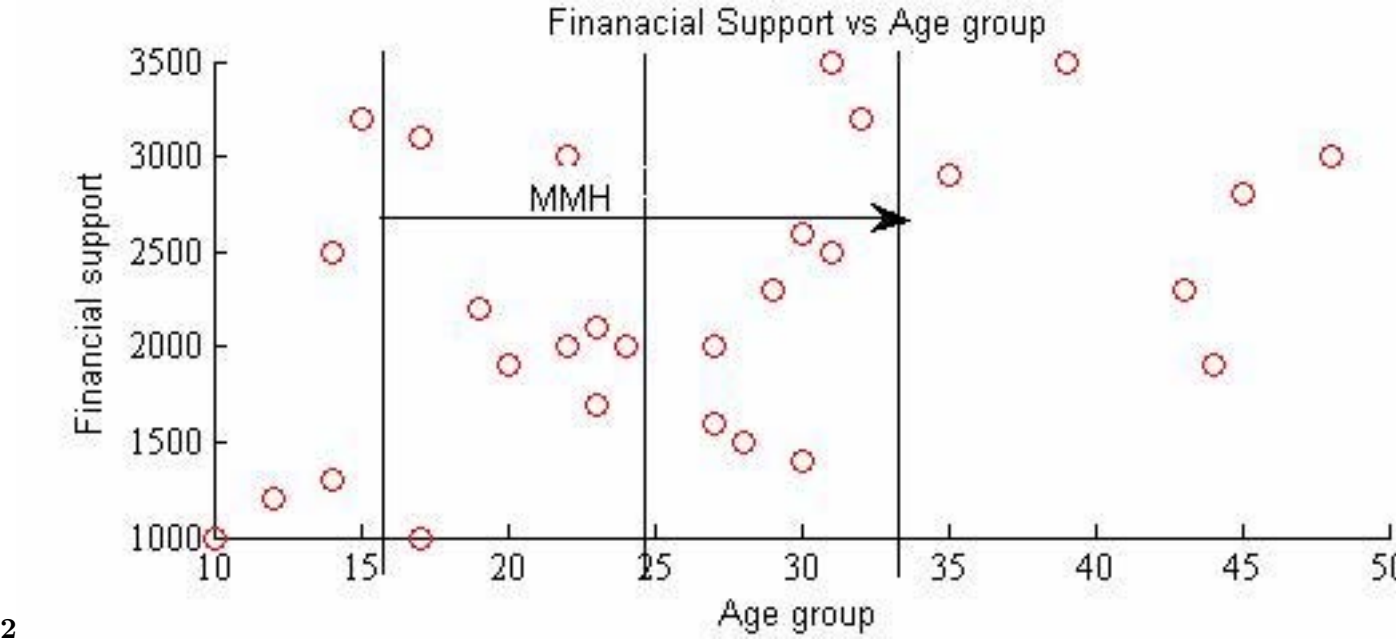


Figure 3: Figure 2 :

221 The total resultant of Bayes Theorem and CART of all data considering financial condition and age group we
222 have got the following table 7:

¹© 2012 Global Journals Inc. (US) 1
²© 2012 Global Journals Inc. (US)

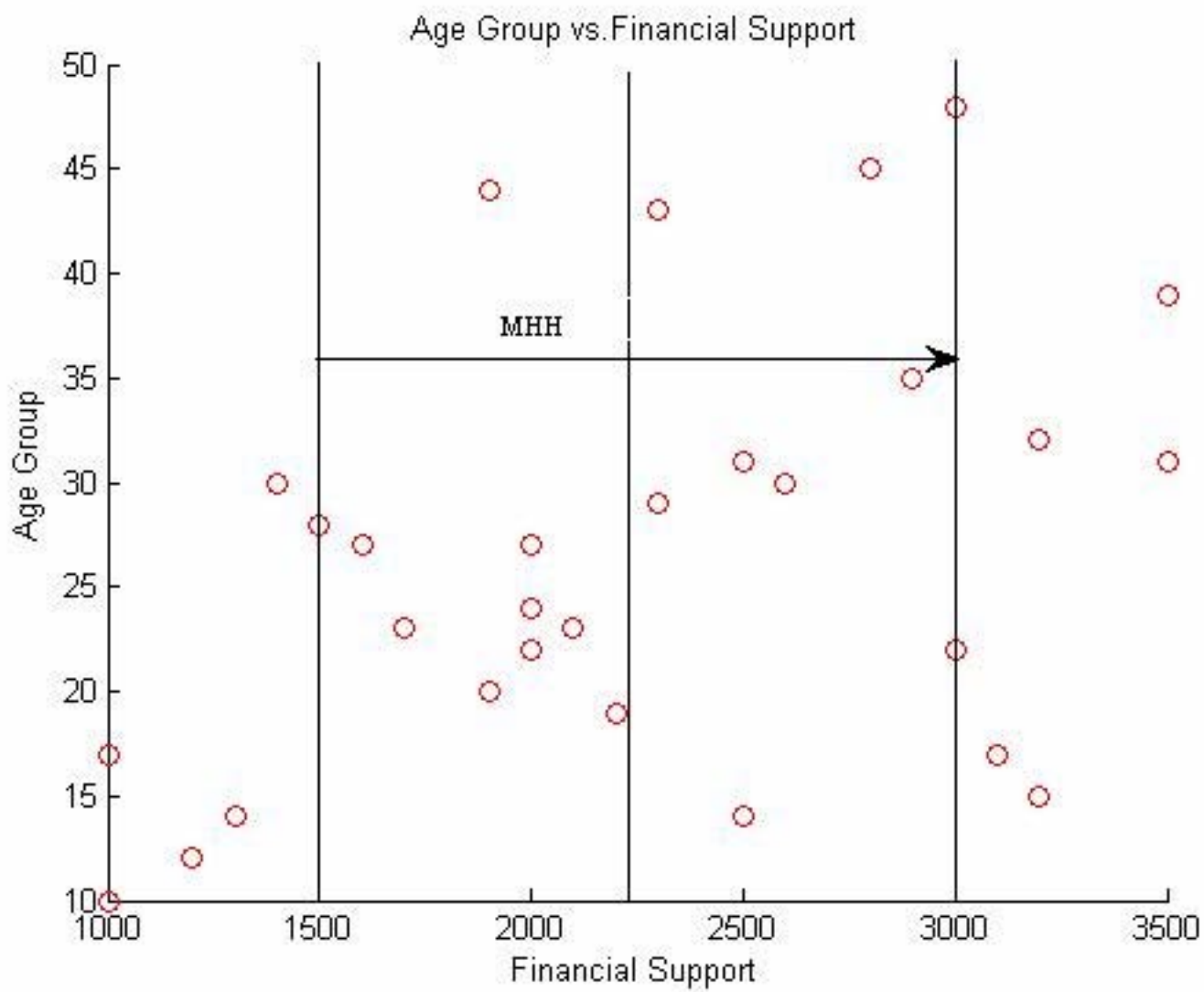


Figure 4:

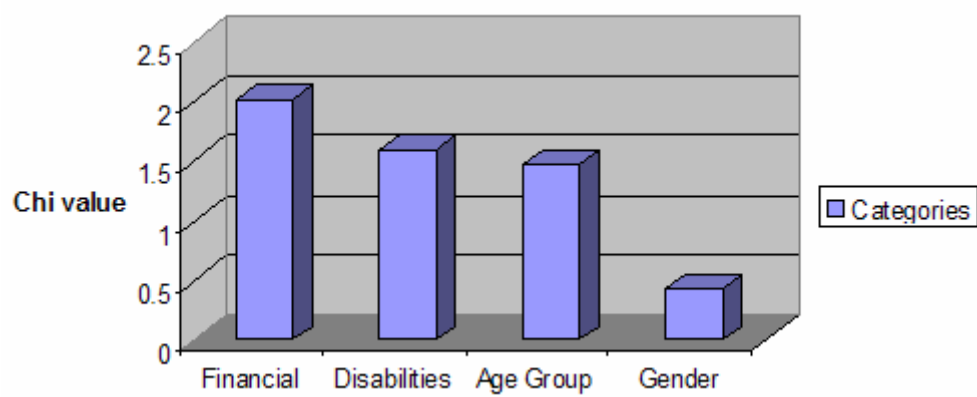
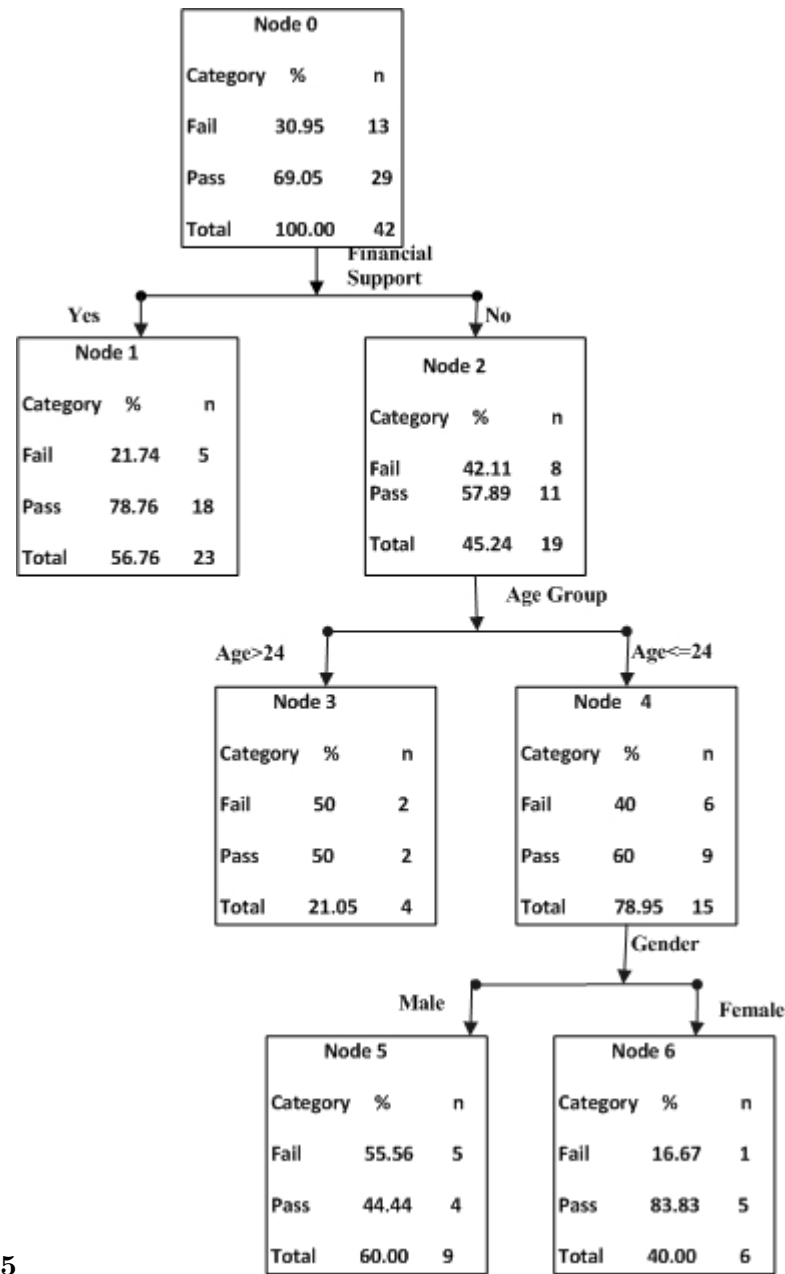


Figure 5: Fig. 4 :



5

Figure 6: Figure 5 :

1

Figure 7: Table 1 :

2

Figure 8: Table 2 : Descriptive statistics (percentage) -Study outcome (42 students)

[illegible]

5

Figure 10: Table 5 :

This study examines the background information from enrolment data that impacts upon the study outcome of Information Systems students at the Department of Computer Science & Engineering. Based on results from feature selection (Figure ?? and Table ??), the CART trees presentation it was found that the most important factors that help separate successful from unsuccessful students are financial support , age group and gender. Demographic data such as gender and age though significantly related to the study outcome, according to the feature selection result.

Based on results from table 5 and 6 by implementing the knowledge of propositional knowledge base and Bayes theorem based on knowledge base to predict the academic success. Better result is found from using Bayes Network.

This study is limited in three main ways that future research can perhaps address. Firstly, this research is based on background information only. Secondly, we used a dichotomous variable for the study outcome with only two categories: pass and fail. Thirdly, from a methodological point of view an alternative to a classification tree should be considered. The prime candidates to be used with this data set are logistic regression and neural networks.

[Artificial Intelligence A modern Approach by Stuart Russell and Peter Norvig] , *Artificial Intelligence A modern Approach by Stuart Russell and Peter Norvig*

[Bean and Metzner ()] ‘A conceptual model of non traditional undergraduate student attrition’. J P Bean , B S Metzner . *Review of Educational Research* 1985. 55 p. .

[Yu et al. ()] ‘A data-mining approach to differentiate predictors of retention’. C H Yu , S Digangi , A Jannasch-Pennell , W Lo , C Kaprolet . *the Proceedings of the Educause Southwest Conference*, (Austin, Texas, USA) 2007.

[Ishitani ()] ‘A longitudinal approach to assessing attrition behavior among first-generation students: Time-varying effects of precollege characteristics’. T T Ishitani . *Research in Higher Education* 2003. 44 (4) p. .

[Pascarella et al. ()] ‘A test and reconceptualization of a theoretical model of college withdrawal in a commuter institution setting’. E T Pascarella , P B Duby , B K Iverson . *Sociology of Education* 1983. 56 p. .

[Boero et al. ()] *An econometric analysis of student withdrawal and progression in post-reform Italian universities. Centro Ricerche Economiche Nord Sud -CRENoS Working Paper*, G Boero , T Laureti , R Naylor . 2005. 2005/04.

[Strayhorn ()] ‘An examination of the impact of first-year seminars on correlates of college student retention’. T L Strayhorn . *Journal of the First-Year Experience & Students in Transition* 2009. 21 (1) p. .

[Rokach and Maimon ()] *Data mining with decision trees -Theory and applications*, L Rokach , O Maimon . 2008. New Jersey: World Scientific Publishing.

[Romero and Ventura ()] ‘Educational data mining: A survey from’. C Romero , S Ventura . *Expert Systems with Applications* 2007. 1995 to 2005. 33 p. .

[Nandeshwar and Chaudhari ()] *Enrollment prediction models using data mining*, A Nandeshwar , S Chaudhari . 2009. 2010.

[Pratt and Skaggs ()] ‘First-generation college students: Are they at greater risk for attrition than their peers?’. P A Pratt , C T Skaggs . *Research in Rural Education* 1989. 6 (1) p. .

[Horstmanshof and Zimitat ()] ‘Future time orientation predicts academic engagement among first-year university students’. L Horstmanshof , C Zimitat . *British Journal of Educational Psychology* 2007. 77 (3) p. .

[Nisbet et al. ()] *Handbook of statistical analysis and data mining applications*, R Nisbet , J Elder , G Miner . 2009. Amsterdam: Elsevier.

[Han and Kamber ()] ‘Investigation of the role of pre-and post admission variables in undergraduate institutional persistence, using a Markov student flow model’. J Han , M Kamber . *Data mining: Concepts and techniques*, (Amsterdam; USA) 2006. 2006. PhD Dissertation. North Carolina State University (2nd ed.)

[Al-Radaideh et al. ()] ‘Mining student data using decision trees’. Q A Al-Radaideh , E M Al-Shawakfa , M I Al-Najjar . *the Proceedings of the 2006 International Arab Conference on Information Technology*, 2006. 2006.

[Md et al. (2012)] ‘New Dropout Prediction for Intelligent System’. Sarwar Md , Linkon Kamal , Chowdhury . *International Journal of Computer Applications* March 2012. (16) p. 42. (Sonia Farhana Nimmy)

[Kember ()] *Open learning courses for adults: A model of student progress*, D Kember . 1995. Englewood Cliffs, NJ: Education Technology.

[Luan and Zhao ()] *Practicing data mining for enrolment management and beyond. New Directions for Institutional Research*, J Luan , C-M Zhao . 2006. 31 p. .

[Tharp ()] ‘Predicting persistence of urban commuter campus students utilizing student background characteristics from enrollment data’. J Tharp . *Community College Journal of Research and Practice* 1998. 22 p. .

- [Dekker et al. ()] *Predicting student drop out: A case study*, G W Dekker , M Pechenizkiy , J M Vleeshouwers .
2009.
- [Simpson ()] ‘Predicting student success in open and distance learning’. O Simpson . *Open Learning*, 2006. 21 p.
- [Noble et al. ()] ‘Predicting successful college experiences: Evi-dence from a first year retention program’. K Noble , N T Flynn , J D Lee , D Hilton . *Journal of College Student Retention: Research, Theory & Practice* 2007. 9 (1) p. .
- [Murtaugh et al. ()] ‘Predicting the retention of university students’. P Murtaugh , L Burns , J Schuster . *Research in Higher Education* 1999. 40 (3) p. .
- [Proceedings of the 2nd International Conference on Educational Data Mining (EDM’09) (2003)] *Proceedings of the 2nd International Conference on Educational Data Mining (EDM’09)*, (the 2nd International Conference on Educational Data Mining (EDM’09)Cordoba, Spain) July 1-3. p. .
- [Reason ()] ‘Student variables that predict retention: Recent research and new developments’. R D Reason . *NASPA Journal* 2003. 40 (4) p. .
- [Ishitani ()] ‘Studying attrition and degree completion behavior among first-generation college students in the United States’. T T Ishitani . *Journal of Higher Education* 2006. 77 (5) p. .
- [Hastie et al. ()] *The elements of statistical learning: Data mining, inference and prediction*, T Hastie , R Tibshirani , J Friedman . 2009. New York: Springer. (2nd ed.)
- [Baker and Yacef ()] ‘The state of educational data mining in 2009: A review and future visions’. R S J D Baker , K Yacef . *Journal of Educational Data Mining* 2009. 1 p. .
- [Siraj and Abdoulha ()] ‘Uncovering hidden information within university’s student enrolment data using data mining’. F Siraj , M A Abdoulha . *MASAUM Journal of Computing* 2009. 1 (2) p. .
- [Jun ()] *Understanding dropout of adult learners in e-learning. PhD Dissertation, the University of*, J Jun . 2005. Georgia, USA.
- [Cortez and Silva ()] ‘Using data mining to predict secondary school student performance’. P Cortez , A Silva . *the Proceedings of 5th Annual Future Business Technology Conference*, (Porto, Portugal) 2008. p. .