

This file has been cleaned of potential threats.

To view the reconstructed contents, please SCROLL DOWN to next page.



FACULTY OF COMPUTERS AND INFORMATION
MENOFIA UNIVERSITY EGYPT

IJCI INTERNATIONAL JOURNAL OF
COMPUTERS AND INFORMATION

Editors:

Prof. Dr. Arabi Keshk

Prof. Dr. Hatem Abdul-Kader

Prof. Dr. Ashraf El Sisi

IJCI

Volume 3 No.1 March 2014

<http://mu.menofia.edu.eg/fci/home/ar>

Hatem6803@yahoo.com

Fainan_nagy1@yahoo.com

International Journal of Computers and Information

Volume 3 No 1

March 2014

Editor-in-chief *Prof. Dr. Arabi Keshk*
Co-Editor-in-chief *Prof. Dr. Hatem Abdul-Kader*
Co-Editor-in-chief *Prof. Dr. Ashraf El Sisi*

Scientific Advisory Editor:

Prof. Dr. Mohiy Mohamed Hadhod	Egypt
Prof. Dr. Nabil Abd El wahed Ismaile	Egypt
Prof. Dr. Fawzy Ali Turkey	Egypt
Prof. Dr. Hany Harb	Egypt
Prof. Dr. Moaed I.M. Dessouky	Egypt
Prof. Dr. Mohamed Kamal Gmal El Deen	Egypt
Prof. Dr. Mahmoud Abd Allah	Egypt
Prof. Nawal Ahmad El Feshawy	Egypt
Prof. Dr. Hegazy Zaher	ISSR
Prof. Dr. Ebrahim Abd El Rahman Farag	ISSR
Prof. Dr. Mohamed Hassan Rasmy	Egypt
Prof. Dr. Hassan Abd El Haleem Yossuf	Egypt
Prof. Dr. Mohamed Abd El Hameed El Esskandarany	Egypt
Prof. Dr. Mohamed Said Ali Ossman	Egypt
Prof. Dr. Abd El Shakoor Sarhan	Egypt
Prof. Dr. Sang. M. Lee.	USA
Prof. Dr. Massimiliano Ferrara	Italy

Journal Secretary

Miss. Fainan Nagy El Sisi

Fainan_nagy1@yahoo.com

TITLE OF THE PAPER (Capital, 14pt Time New Roman, Bold, Centered)

Author Name (12pt, Time New Roman, Centered)

Author address (10pt, Time New Roman, Centered)

E_mail@server.xxx (10pt, Time New Roman, Centered)

Abstract: this the Sample format of your full paper. Use Word for windows (Microsoft) or equivalent word processor with exactly the same "printing result" by tuning 2cm from right and 2cm from left in the Microsoft word package, or equivalently by keeping 2.5cm, real distance, from right and from left. Use single space. Use one column format paper after the keywords 11Pt Time New Roman For Abstract, and KEYWORDS USE Italic.

Keywords: - Leave one blank line after the Abstract and write you Keywords (6- 11 words).

1. Introduction

As you can see for the title of paper you must use 14Pt, centered, bold, Time New Roman. Leave one Blank line and then type Authors' Name (Capitalize each word, 12Pt, Time New Roman, centered). Address(in 12 Pt Time New Roman, Centered). Then you must type your email address e_mail@server.xxx and your website address <http://www.yourwebaddress.xxx.xx> both in 10Pt Time New Roman, Centered).

The heading of each section should be printed in small 14Pt, left justified, bold, Time New Roman. You must use numbers 1, 2, 3 ... for the section numbering and not Latin numbering (1.11, 111 ...).

2. Problem Formulation

Please leave two blank lines between successive sections as here.

Mathematical Equations must be numbered as follows: (1), (2), ..., (990.nd not (1.1), (1.2)..., etc. depending on your various sections.

3. Subsection

When including a sub-subsection you must use for its heading, small letters, 12Pt, left justified, bold Time New Roman as here.

3.1 sub-subsection

When including a sub-subsection you must use for its heading, small letters, 11Pt, left justified, bold Time New Roman as here.

4. Problem Solution

Figures and Tables should be numbered. 'S follow: Figure1, Figure2... etc. Table1, Table2 ... etc. if your paper deviates significantly from these specifications. Our Publishing house may not be able to include your paper in the proceedings. When citing references in the text of the abstract, type the corresponding number in square brackets as shown at the end of this sentence (1)

5. Conclusion

Please follow our instruction faithfully, otherwise you resubmit your full paper. This will enable us to maintain uniformity in the journal proceedings. Thank you for your cooperation.

6. References

- [1] X1. Author, Title of Paper, International Journal of Computers and information, Vol. X, NO. X, 20XX, pp. XX-XX.**
- [2] X2. Author, Title of Paper, Title of the book, Publishing House,20XX.**

Table of Contents

Paper Title	PN.
Evaluation of Differential Evolution and Particle Swarm Optimization Algorithms at Training of Neural Network for prediction Abdul Salam, M.E, H.M.Abdulkader, W.F. Abdul Wahed	2
A Novel Rhetorical Structure Approach for Classifying Arabic Security Documents H. Mathkour	15
Visual Inspection of Ceramic Tiles Surfaces Using Statistical Features and LVQ of Artificial Neural Networks I. El-Henawy, S. Elmougy, A. El-Azab	28
Evaluation of MUVES: Needs and Results <i>S. M. Abd El-razek, H. M. EL-bakry, W.F. Abd El-wahd</i>	42
An Efficient Technique For SQL Injection Detection And Prevention E. M. SAFWAT, H. MAHGOUB, A.EL-SISI, A. KESHK	57
3D Object Categorization Using Spin-Images with MPI Parallel Implementation A. Eleliemy, D. Hegazy, W. Elkilani	75

Evaluation of Differential Evolution and Particle Swarm Optimization Algorithms at Training of Neural Network for prediction

Abdul Salam, M.E

Teaching Assistant at Higher
Technological Institute (H.T.I)
10th of Ramadan City, Egypt
mustafa.abdo@ymail.com

H.M.Abdulkader

Faculty of Computer and
Information, Menoufiya University,
Egypt
hatem6803@yahoo.com

W.F. Abdul Wahed,

Head of Operations Research and
Decision Support Systems, Faculty of
Computers and Information, Menoufiya
University, Egypt

Abstract: This paper presents the comparison of two metaheuristic approaches: Differential Evolution (DE) and Particle Swarm Optimization (PSO) in the training of feed-forward neural network to predict the daily stock prices. Stock market prediction is the act of trying to determine the future value of a company stock or other financial instrument traded on a financial exchange. The successful prediction of a stock's future price could yield significant profit. The feasibility, effectiveness and generic nature of both DE and PSO approaches investigated are exemplarily demonstrated. Comparisons were made between the two approaches in terms of the prediction accuracy and convergence characteristics. The proposed model is based on the study of historical data, technical indicators and the application of Neural Networks trained with DE and PSO algorithms. Results presented in this paper show the potential of both algorithms applications for the decision making in the stock markets, but DE gives better accuracy compared with PSO.

Keywords:- Evolutionary algorithms, Differential evolution, Particle swarm optimization, feed-forward neural network, technical indicators, stock prediction.

1. Introduction

STOCK price prediction has been at focus for years since it can yield significant profits. Predicting the stock market is not a simple task, mainly as a consequence of the close to random-walk behaviour of a stock time series. Fundamental and technical analysis was the first two methods used to forecast stock prices. Neural networks are the most commonly used technique [1]. The role of artificial neural networks in the present world applications is gradually increasing and faster algorithms are being developed for training neural networks [2]. In general, back-propagation is a method used for training neural networks. Gradient descent, conjugate gradient descent, resilient, BFGS quasi-Newton, one-step secant, Levenberg-Marquardt and Bayesian regularization are all different forms of the back-propagation training algorithm. For all these algorithms storage and computational requirements are different, some of these are good for pattern recognition and others for function approximation but they have drawbacks in one way or other, like neural network size and their associated storage requirements. Certain training algorithms are suitable for some type of applications only, for example an algorithm that performs well for pattern recognition may not for classification

problems and vice versa, in addition some cannot cater for high accuracy/performance. It is difficult to find a particular training algorithm that is the best for all applications under all conditions all the time [3]. The perceived advantages of evolution strategies as optimization methods motivated the authors to consider such stochastic methods in the context of training artificial neural networks and optimizing the structure of the networks [4]. A survey and overview of evolutionary algorithms in evolving artificial neural networks can be found in [5].

Differential evolution (DE) is introduced by Kenneth Price and Rainer Storn in 1995. DE algorithm is like genetic algorithms using similar operators; crossover, mutation and selection. DE can find the true global minimum regardless of the initial parameter values. The main difference in constructing better solutions is that genetic algorithms rely on crossover while DE relies on mutation operation. DE is successfully applied to many artificial and real optimization problems and applications, such as aerodynamic shape optimization [16], automated mirror design [3], optimization of radial active magnetic bearings [18], and mechanical engineering design [10]. A differential evolution based neural network training algorithm, introduced in [5,10,23]. PSO is proposed algorithm by James Kennedy and Russell Eberhart in 1995, motivated by social behavior of organisms such as bird flocking and fish schooling [6]. The main difference of particle swarm optimization concept from the evolutionary computing is that flying potential solutions through hyperspace are accelerating toward "better" solutions, while in evolutionary computation schemes operate directly on potential solutions which are represented as locations in hyperspace [9]. Neural networks are used in combination with PSO in many applications, like neural network control for nonlinear processes in [10], feedforward neural network training in [11] – [17]. PSO algorithm is used in prediction and forecasting in many applications, like prediction of chaotic systems in [23], electric load forecasting in [25],[26], time series prediction in [28]-[30] and stock market decision making in [31]-[33].

The main Purpose of this paper is to propose of using two modern Artificial intelligence technique in neural network training. Stock predication is selected as an application to evaluate the efficiency of the both proposal training Algorithms for its complexity and sensitivity.

The paper is organized as follows: Section 2 presents the Differential Evolution algorithm; Section 3 presents Particle swarm optimization algorithm; Section 4 is devoted for the proposed system and implementation of Differential Evolution and particle swarm optimization algorithms in stock prediction; In Section 5 the results are discussed. The main conclusions of the work are presented in Section 6.

2. Differential evolution training algorithm

Price and Storn developed DE to be a reliable and versatile function optimizer. The first written publication on DE appeared as a technical report in 1995 (Price and Storn 1995). Like nearly all EAs, DE is a population-based optimizer that attacks the starting point problem by sampling the objective function at multiple, randomly chosen initial points. [19]. DE algorithm like genetic algorithms using similar operators; crossover, mutation and selection. DE has three advantages; finding the true global minimum regardless of the initial parameter values, fast convergence, and using few control parameters. The main difference in constructing better solutions is that genetic algorithms rely on crossover while DE relies on mutation operation. This main operation is based on the differences of randomly sampled pairs of solutions in the population. The algorithm uses mutation operation as a search mechanism and selection operation to direct the search toward the prospective regions in the search space. The DE algorithm also uses a non-uniform crossover that can take child vector parameters from one parent more often than it does from others. By using the components of the existing population members to construct trial vectors, the recombination (crossover) operator efficiently shuffles information about successful combinations, enabling the search for a better solution space [20]. The DE algorithm is shown in figure 1.

2.1 Population Structure

The current population, symbolized by P_x , is composed of those vectors, $x_{i,g}$, that have already been found to be acceptable either as initial points, or by comparison with other vectors:

$$Px,g = (x_{i,g}), i=0,1,...,Np-1, g=0,1,...,gmax, x_{i,g}=(x_{j,i,g}), j=0,1,...,D-1 \quad (1)$$

Once initialized, DE mutates randomly chosen vectors to produce an intermediary population, $P_{v,g}$, of N_p mutant vectors, $V_{i,g}$:

$$\begin{aligned} P_{v,g} &= (V_{i,g}), i=0,1,...,Np-1, g=0,1,...,gmax, \\ V_{i,g} &= (V_{j,i,g}), j=0,1...D-1 \end{aligned} \quad (2)$$

Each vector in the current population is then recombined with a mutant to produce a trial population, P_u , of N_p trial vectors, $u_{i,g}$:

$$\begin{aligned} P_{u,g} &= (u_{i,g}), i=0,1,...,Np-1, g=0,1,...,gmax, \\ u_{i,g} &= (u_{j,i,g}), j=0,1,...,D-1 \end{aligned} \quad (3)$$

During recombination, trial vectors overwrite the mutant population, so a single array can hold both populations.

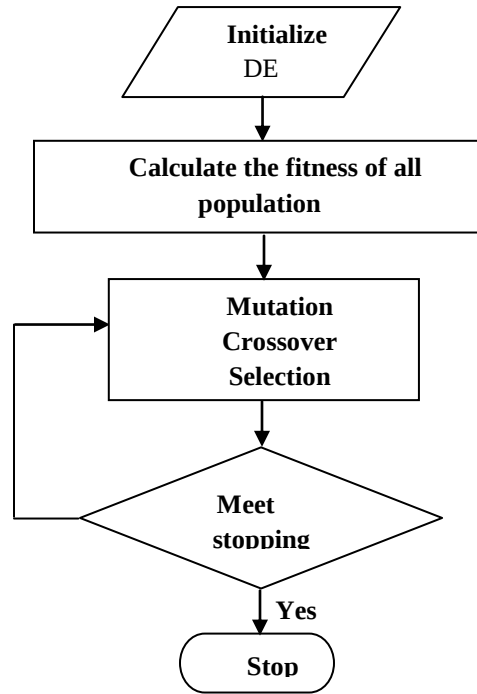


Figure 1 the DE algorithm

2.2 Initialization

Before the population can be initialized, both upper and lower bounds for each parameter must be specified. These 2D values can be collected into two, D-dimensional initialization vectors, b_L and b_U . Once initialization bounds have been specified, a random number generator assigns each parameter of every vector a value from within the prescribed range. For example, the initial value ($g = 0$) of the j th parameter of the i th vector is

$$x_{j,i,0} = \text{rand}_j(0,1) \cdot (b_{j,U} - b_{j,L}) + b_{j,L}. \quad (4)$$

2.3 Mutation

Once initialized, DE mutates and recombines the population to produce a population of N_p trial vectors. In particular, differential mutation adds a scaled, randomly sampled, vector difference to a third vector.

$$V_{i,g} = x_{r0,g} + F \cdot (x_{r1,g} - x_{r2,g}) \quad (5)$$

The scale factor, $F \in (0,1+)$, is a positive real number that controls the rate at which the population evolves. While there is no upper limit on F , effective values are seldom greater than 1.0.

2.4 Crossover

To complement the differential mutation search strategy, DE also employs uniform crossover. Sometimes referred to as discrete recombination, (dual) crossover builds trial vectors out of parameter values that have been copied from two different vectors. In particular, DE crosses each vector with a mutant vector:

$$u_{i,g} = (u_{j,i,g}) = \begin{cases} v_{i,i,g} & \text{if } \text{rand}_i(0,1) \leq Cr \text{ or} \\ j = j_{\text{rand}} & \\ x_{i,i,g} & \text{otherwise} \end{cases} \quad (6)$$

The crossover probability, $Cr \in [0,1]$, is a user-defined value that controls the fraction of parameter values that are copied from the mutant.

2.5 Selection

If the trial vector, $u_{i,g}$, has an equal or lower objective function value than that of its target vector, $x_{i,g}$, it replaces the target vector in the next generation; otherwise, the target retains its place in the population for at least one more generation

$$x_{i,g+1} = \begin{cases} u_{i,g} & \text{if } f(u_{i,g}) \leq f(x_{i,g}) \\ x_{i,g} & \text{otherwise} \end{cases} \quad (7)$$

Once the new population is installed, the process of mutation, recombination and selection is repeated until the optimum is located, or a prespecified termination criterion is satisfied, e.g., the number of generations reaches a preset maximum, g_{max} . [18]

One possibility could be a hybrid of traditional optimization methods and evolutionary algorithms as studied in [2, 4, 11].

3. PARTICLE SWARM OPTIMIZATION ALGORITHM

PSO is a relatively recent heuristic search method which is derived from the behavior of social groups like bird flocks or fish swarms. PSO moves from a set of points to another set of points in a single iteration with likely improvement using a combination of deterministic and probabilistic rules. The PSO has been popular in academia and industry, mainly because of its

intuitiveness, ease of implementation, and the ability to effectively solve highly nonlinear, mixed integer optimization problems that are typical of complex engineering systems. Although the “survival of the fittest” principle is not used in PSO, it is usually considered as an evolutionary algorithm. Optimization is achieved by giving each individual in the search space a memory for its previous successes, information about successes of a social group and providing a way to incorporate this knowledge into the movement of the individual. Therefore, each individual (called particle) is characterized by its position $\bar{x}_i$, its velocity $\bar{v}_i$, its personal best position $\bar{p}_i$ and its neighborhood best position $\bar{p}_g$.

The elements of the velocity vector for particle i are updated as

$$v_{ij} \leftarrow \omega v_{ij} + c_1 q(x_{ij}^{pb} - x_{ij}) + c_2 r(x_j^{sb} - x_{ij}), j=1,..,n \quad (8)$$

Where w is the inertia weight, x_i^{pb} is the best variable vector encountered so far by particle i , and x^{sb} is the swarm best vector, i.e. the best variable vector found by any particle in the swarm, so far. c_1 and c_2 are constants, and q and r are random numbers in the range $[0, 1]$. Once the velocities have been updated, the variable vector of particle i is modified according to

$$x_{ij} \leftarrow x_{ij} + v_{ij}, j=1,...,n. \quad (9)$$

The cycle of evaluation followed by updates of velocities and positions (and possible update of x_i^{pb} and x^{sb}) is then repeated until a satisfactory solution has been found. PSO algorithm is shown in figure 2.

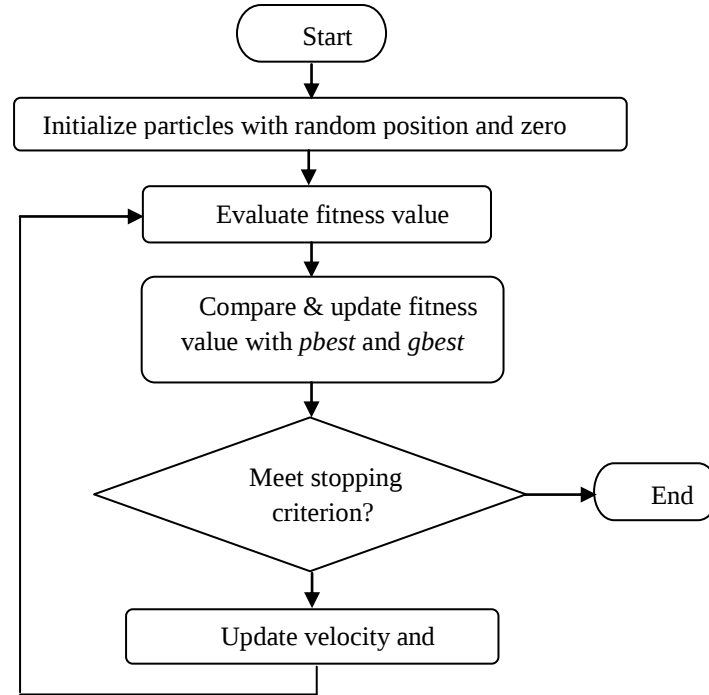


Figure 2 PSO algorithm

4. THE PROPOSED MODEL

The proposed methodology is to train multilayer feed forward neural network with Differential Evolution (DE) algorithm and also Particle Swarm Optimization algorithm to be used in the prediction of daily stock prices. The proposed model is based on the study of historical data, technical indicators and the application of Neural Networks trained with DE and PSO algorithms. Neural network architecture contains one input layer with six inputs neurons represent the historical data and derived technical indicators, one hidden layers and single output layer as shown in Figure 3.

The two algorithms were tested for many companies which cover different stock sectors, like Drug Manufacturers, Industries, Utilities, Communications, life science and Automotives. These companies are Acadia Pharmaceuticals Inc. (ACAD), Shiloh Industries Inc. (SHLO), FiberTower Corporation (FTWR), Hayes Lemmerz International Inc. (HAYZ), Strategic Internet (SIIL.OB), Caliper Life Sciences, Inc. (CALP) and Ford.

Five technical indicators are calculated from the raw datasets for neural networks inputs: Relative Strength Index (RSI): A technical momentum indicator that compares the magnitude of recent gains to recent losses in an attempt to determine overbought and oversold conditions of an asset. The formula for computing the Relative Strength Index is as follows.

$$RSI = 100 - [100 / (1 + RS)] \quad (10)$$

Where $RS = \text{Avg. of } x \text{ days' up closes} / \text{Average of } x \text{ days' down closes}$.

Money Flow Index (MFI): This one measures the strength of money in and out of a security. The formula for MFI is as follows.

$$\text{Money Flow (MF)} = \text{Typical Price} * \text{Volume.} \quad (11)$$

$$\text{Money Ratio (MR)} = (\text{Positive MF} / \text{Negative MF}). \quad (12)$$

$$\text{MFI} = 100 - (100 / (1 + \text{MR})). \quad (13)$$

Exponential Moving Average (EMA): This indicator returns the exponential moving average of a field over a given period of time. EMA formula is as follows.

$$\text{EMA} = [\alpha * \text{Today's Close}] + [1 - \alpha * \text{Yesterday's EMA}]. \quad (14)$$

Stochastic Oscillator (SO): The stochastic oscillator defined as a measure of the difference between the current closing price of a security and its lowest low price, relative to its highest high price for a given period of time. The formula for this computation is as follows.

$$\%K = [(\text{Close price} - \text{Lowest price}) / (\text{Highest Price} - \text{Lowest Price})] * 100 \quad (15)$$

Moving Average Convergence/Divergence (MACD): This function calculates difference between a short and a long term moving average for a field. The formulas for calculating MACD and its signal as follows.

$$\text{MACD} = [0.075 * \text{EMA of Closing prices}] - [0.15 * \text{EMA of closing prices}] \quad (16)$$

$$\text{Signal Line} = 0.2 * \text{EMA of MACD} \quad (17)$$

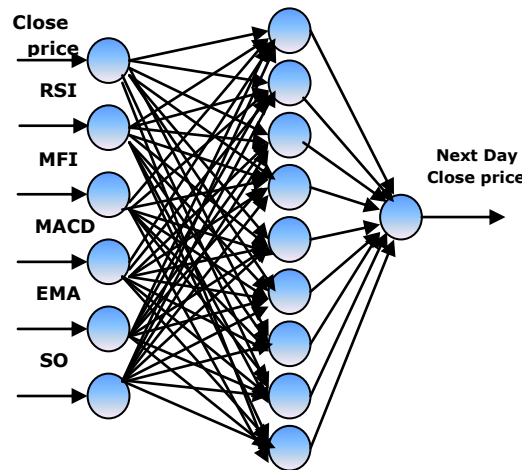


Figure3. Architecture of neural network of proposed model.

5. Results and Discussion

Neural networks are trained and tested with datasets from September 2004 to September 2007. All datasets are available on <http://finance.yahoo.com> web site. Datasets are divided into training part (70%) and testing part (30%).

The used software is Matlab and Microsoft excel.

Figures (4-10) outlines the application of different training algorithms at different data sets with different sectors of the market. The used data sets present different market trends. In figures (4, 7) which present results of the two algorithms on two different sectors which are Pharmaceuticals and utilities sectors; one can remark that the predicted curve using both DE and PSO algorithms give a good accuracy with a little advance to DE algorithm and the datasets are not fluctuated.

Figures (5, 8, 9) outline the application of the two algorithms on a different market sectors. From figures one can remark the enhancement in the error rate achieved by the DE algorithm.

Figure 6, 10 outlines different time series with changing trends. The DE can easily cope up with the fluctuation existing in the time series better than PSO algorithm.

performance as the weight sum of two factors: the mean squared error and the mean squared weights and biases; mean square error (MSE) and mean absolute error (MAE) performance functions. It can be remarked that the DE always gives an advance over the PSO algorithms in all performance functions and in all trends and sectors. DE performs better than PSO especially in cases with fluctuations in the time series function.

Table 1. outlines Mean squared error with regularization performance function (MSEREG) which measures network

Error Company	MSEREG		MSE		MAE	
	DE	PSO	DE	PSO	DE	PSO
ACADIA	0.07	1.50	0.49	1.46	0.46	0.90
SHLO	1.04	1.43	1.01	1.39	0.80	0.90
FTWR	0.17	0.48	0.16	0.30	0.30	0.49
HYZ	0.06	0.19	0.00	0.11	0.17	0.20
NET	0.71	2.27	0.70	1.94	0.72	1.14
CALPER	0.00	0.16	0.02	0.09	0.12	0.20
FORD	0.34	2.10	0.22	2.09	0.39	1.30
Table (1) Error functions for the two algorithms.						

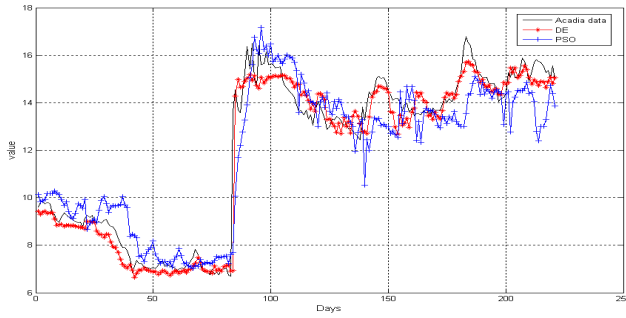


Figure 4 Results for Acadia Pharmaceuticals Company.

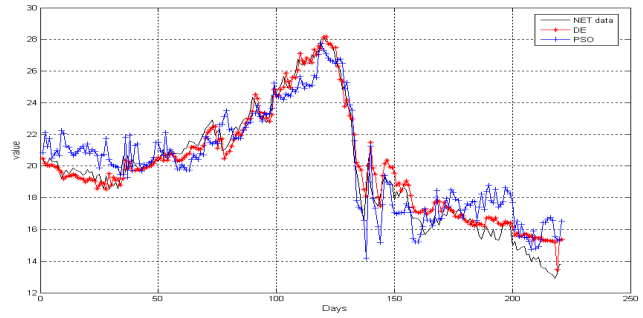


Figure 8 Results for Strategic Internet (SIIL.OB) Company.

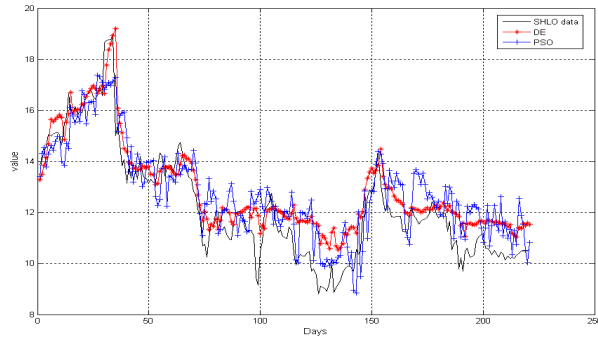


Figure 5 Results for Shiloh Industries Company.

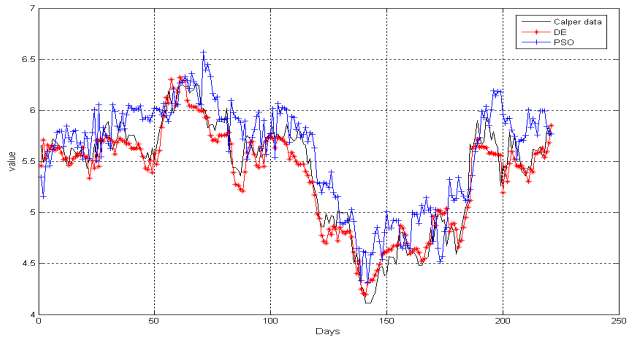


Figure 9 Results for Caliper Life Sciences Company.

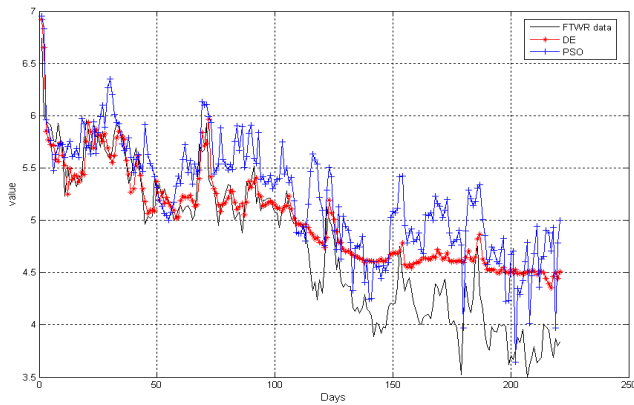


Figure 6 Results for Fiber Tower Company.

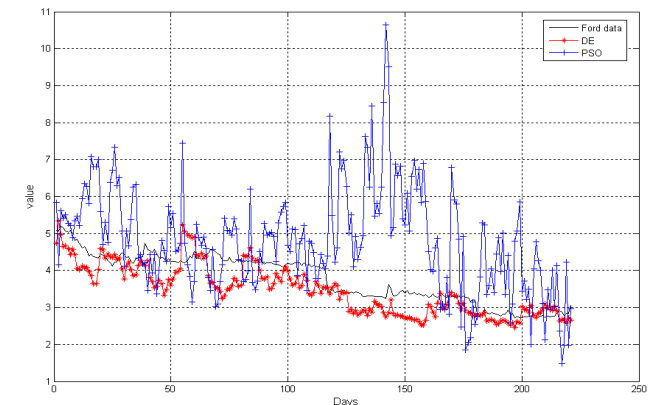


Figure 10 Results for Ford Company.

6. CONCLUSIONS

In this paper, Differential Evolution (DE) and Particle Swarm Optimization (PSO) algorithms are applied in the training and testing of feed-forward neural network for stock price prediction. The simulation results show the potential of both two algorithms. Both algorithms Convergence

to a global minimum can be expected. Both algorithms can avoid local minima problem which all gradient descending methods fall on it. Easy tuning of both algorithms parameters. Optimum found by both algorithms is never worse than the initial optimum found by a gradient based method. But DE converges to global minimum faster than PSO algorithm. DE algorithm gives better accuracy than PSO algorithm especially in fluctuated time series.

6. References

- [1] Olivier coupelon, Blaise Pascal University: Neural network modeling for stock movement prediction, state of art. 2007
- [2] Howard Demuth Mark Beale Martin Hagan: Neural Network Toolbox for Use with MATLAB® 2006.
- [3] Tejen Su, Jyunwei Jhang and Chengchih Hou: A hybrid artificial neural networks and particle swarm optimization for function approximation; ICIC International 2008 ISSN 1349-4198.
- [4] B. Al-kazemi and C.K. Mohan. Training feedforward neural networks using multiphase particle swarm optimization. In Neural Information Processing, 2002. ICONIP '02. Proceedings of the 9th International Conference on, pages 2615 – 2619 vol.5, 2002.
- [5] Yao,X.: Evolving artificial neural networks, Proceedings of the IEEE, 87(9) (1999),1423–1447.
- [6] D. N. Wilke. Analysis of the particle swarm optimization algorithm, Master's Dissertation, University of Pretoria, 2005.
- [7] Jovita Nenortaite and Rimvydas Simutis: Application of Particle Swarm Optimization Algorithm to Stocks' Trading System.2004.
- [8] Khalil A.S.: An Investigation into Optimization Strategies of Genetic Algorithms and Swarm Intelligence.. Artificial Life (2001).
- [9] Kennedy J., Spears W.M.: Matching Algorithms to Problems: An Experimental Test of the Particle Swarm and Some Genetic Algorithms on the Multimodal Problem
Generator.Processes.<http://www.aic.nrl.navy.mil/%7Espears/papers/wcci98.pdf>.
Current as of December 15th, 2003.
- [10] Ying Song, Zengqiang Chen, and Zhuzhi Yuan. New chaotic pso-based neural network predictive control for nonlinear process. IEEE Transactions on Neural Networks, 18:595 –601, 2007.
- [11] M. Carvalho and T.B. Ludermir. Particle swarm optimization of feed-forward neural networks with weight decay. In Hybrid Intelligent Systems, 2006. HIS '06. Sixth International Conference on, pages 5 – 5, 2006.
- [12] Zhang Chunkai, Li Yu, and Shao Huihe. A new evolved artificial neural network and its application. In Intelligent Control and Automation, 2000. Proceedings of the 3rd World Congress on, pages 1065 – 1068 vol.2, 2000.

- [13] Hong-Bo Liu, Yi-Yuan Tang, Jun Meng, and Ye Ji. Neural networks learning using vbest model particle swarm optimisation. In Proceedings of 2004 International Conference on Machine Learning and Cybernetics, pages 3157 – 3159 vol.5, 2004.
- [14] R. Mendes, P. Cortez, M. Rocha, and J. Neves. Particle swarms for feedforward neural network training. In Neural Networks, 2002. IJCNN '02. Proceedings of the 2002 International Joint Conference on, pages 1895 – 1899, 2002.
- [15] Cui-Ru Wang, Chun-Lei Zhou, and Jian-Wei Ma. An improved artificial fish-swarm algorithm and its application in feed-forward neural networks. In Machine Learning and Cybernetics, 2005. Proceedings of 2005 International Conference on, pages 2890 – 2894 Vol.5, 2005.
- [16] Chunkai Zhang, Huihe Shao, and Yu Li. Particle swarm optimisation for evolving artificial neural network. In 2000 IEEE International Conference on Systems, Man, and Cybernetics, pages 2487 – 2490 vol.4, 2000.
- [17] Fuqing Zhao, Zongyi Ren, Dongmei Yu, and Yahong Yang. Application of an improved particle swarm optimization algorithm for neural network training. In Neural Networks and Brain, 2005. ICNN&B '05. International Conference on, pages 1693 – 1698, 2005.
- [18] Chia-Feng Juang. A hybrid of genetic algorithm and particle swarm optimization for recurrent network design. IEEE Transactions on Systems, Man and Cybernetics, Part B, 34:997– 1006, 2004.
- [19] Chia-Feng Juang and Yuan-Chang Liou. On the hybrid of genetic algorithm and particle swarm optimization for evolving recurrent neural network. In Neural Networks, 2004. Proceedings. 2004 IEEE International Joint Conference on, pages 2285 – 2289 vol.3, 2004.
- [20] H.A. Firpi and E.D. Goodman. Designing templates for cellular neural networks using particle swarm optimization. In Applied Imagery Pattern Recognition Workshop, 2004. Proceedings. 33rd, pages 119 – 123, 2004.
- [21] Yuehui Chen, Jiwen Dong, Bo Yang, and Yong Zhang. A local linear wavelet neural network. In Intelligent Control and Automation, 2004. WCICA 2004. Fifth World Congress on, pages 1954 – 1957 Vol.3, 2004.
- [22] P.D. Reynolds, R.W. Duren, M.L. Trumbo, and II Marks R.J. Fpga. Implementation of particle swarm optimization for inversion of large neural networks. In Proceedings 2005 IEEE Swarm Intelligence Symposium, 2005. SIS 2005, pages 389 – 392, 2005.
- [23] F.A. Guerra and L.D.S. Coelho. Radial basis neural network learning based on particle swarm optimization to multistep prediction of chaotic lorenz's system. In Hybrid Intelligent Systems, 2005. Fifth International Conference on, page 3 pp., 2005.
- [24] Changhui Deng, XinJiang Wei, and LianXi Guo. Application of neural network based on pso algorithm in prediction model for dissolved oxygen in fishpond. In

- Intelligent Control and Automation, 2006. WCICA 2006. The Sixth World Congress on, pages 9401 – 9405, 2006.
- [25] Changyin Sun and Dengcai Gong. Support vector machines with pso algorithm for short-term load forecasting. In 2006. ICNSC '06. Proceedings of the 2006 IEEE International Conference on Networking, Sensing and Control, pages 676– 680, 2006.
- [26] Wei Sun, Ying-Xia Zhang, and Fang-Tao Li. The neural network model based on pso for short-term load forecasting. In Machine Learning and Cybernetics, 2006 International Conference on, pages 3069 – 3072, 2006.
- [27] Jinchun Peng, Yaobin Chen, and R. Eberhart. Battery pack state of charge estimator design using computational intelligence approaches. In Battery Conference on Applications and Advances, 2000. The Fifteenth Annual, pages 173 – 177, 2000.
- [28] X. Cai, N. Zhang, G.K. Venayagamoorthy, and II Wunsch D.C. Time series prediction with recurrent neural networks using a hybrid pso-ea algorithm. In Neural Networks, 2004. Proceedings. 2004 IEEE International Joint Conference on, pages 1647 – 1652 vol.2, 2004.
- [29] Y. Yang, R.S. Chen, Z.B. Ye, and Z. Liu. Fdtd. time series extrapolation by the least squares supports vector machine method with the particle swarm optimization technique. In Microwave Conference Proceedings, 2005. APMC 2005. Asia-Pacific Conference Proceedings, page 3 pp., 2005.
- [30] Chunkai Zhang and Hong Hu. Using pso algorithm to evolve an optimum input subset for a svm in time series forecasting. In 2005 IEEE International Conference on Systems, Man and Cybernetics, pages 3793 – 3796 Vol. 4, 2005.
- [31] Jovita Nenortaitė. A particle swarm optimization approach in the construction of decision-making model. ISSN 1392 -124x information technology and control, 2007, vol.36, no.1a
- [32] N.G. Pavlidis, D. Tasoulis, M.N. Vrahatis. Financial Forecasting Through Unsupervised Clustering and Evolutionary Trained Neural Networks. 2003 Congress on Evolutionary Computation, Canberra Australia, 2003.
- [33] J. Nenortaite, R. Simutis. Stocks' Trading System Based on the Particle Swarm Optimization Algorithm. Lecture Notes on Computer Science, 2004, 3039, 843-850.

A Novel Rhetorical Structure Approach for Classifying Arabic Security Documents

Hassan Mathkour

Department of Computer Science

King Saud University

Riyadh 11653

Saudi Arabia

binmathkour@yahoo.com

mathkour@ccis.ksu.edu.sa

Abstract: Security Documents classification is aimed at securing documents from being illegally disclosed. Classifying a portion of a document as a 'secret' depends on the type of effect its disclosure will have in an organization. In this respect, Information is classified according to their critical semantic (i.e. its context or value and intended uses or audience at particular time or situation). Understanding the semantic of a document is not an easy task. The rhetorical structure theory (RST) is one of the leading theories that have been applied successfully in text processing and understanding. In this paper, we will describe a novel approach to automatically classify Arabic Security documents using RST.

Keywords:- RST, Information classification, Arabic security documents, automatic document classification.

3. Introduction

Profitability of organizations is ultimately dependent on the effectiveness with which they exchange, process, control, manage and, more importantly, protect their data and information so that only authorized persons access them. All these processes require that the right information be made available to the authorized persons at the right place and at the right time [9]. To this end, one needs to provide a well-defined methodology for specifying the criteria that govern the security classification. In addition, it necessary to re-evaluate the value and importance of the information being classified to reassign the security as this changes with time depending on the organization. This will require defining appropriate procedures and protection requirements for the security reassignment. Not all documents are of the same importance and so not all information in those documents requires the same degree of protection. This implies that different portions of information are assigned different security classifications. The assignment follows different schemes depending on the nature of the organization.

The first step in information classification is to identify a senior member of management as the authorized person of the particular information to be classified. The second step is to develop a classification policy. The policy should describe the different classification labels, define the criteria for assigning a particular label to information, and then list the required security controls for each classification [9].

Volubility of the information in the document and their timeliness are among the factors that influence the assignment of document security classification. Laws, national security issues and other regulatory requirements are important considerations when classifying documents. Furthermore, in classifying documents, the security requirements of specific categories of documents, the various processing stages of documents, such as draft and final, and the contents and structure of documents must be clearly articulated and specified. A review of the classification of particular information should be done periodically to ensure that its classification is still appropriate, and also to ensure the security controls required by the classification are in place [7, 20, 21].

Documents are classified according to their importance [1], or their potential effect on a specific organization. Classified information is sensitive information to which access by particular classes of people is restricted by either law or regulation. A formal security clearance is required to handle classified documents or access classified data [20, 21]. The common information security classification labels that are used by the business sector are public, sensitive, private, and confidential. Government and its agencies adopt classification labels such as Classified, Unclassified, Sensitive but Unclassified, Restricted, Confidential, Secret, Top Secret and their non-English equivalents [10, 21].

Top secret- The compromise of this information or material would likely cause a tremendous effect on the organization. Secret- The compromise of this information or material would likely cause a big effect (but less than the above class) on the organization and should not be disclosed by anyone.

Confidential- Indicates that the information is private information which should not be disclosed by a person other than to the intended one.

The term "unclassified" was not formally a classification but could be used to indicate positively that information or material did not carry a classification.

Companies have documents that carry private information. They need to label them with one of the labels described above to inform the recipient about their importance. Classifying the information in a certain document depends mainly on its semantic (i.e. its context or value and intended uses or audience at particular time or situation). An example is that part of the document that contains the salaries of the executive board members. In some cases, this information may be considered as secret in a document of the financial analysis of the company. This same information may be considered as unclassified if mentioned in the context of the richest people in the city [7, 20, 21].

It is most likely that the classification is done manually by some people who follow certain guidelines. This process may take long time when the document is huge –hundreds of pages. However, performing the classification automatically would make the process faster and easier. The idea of automating the process has the challenge of understanding the semantic of the information. In this paper, we propose a novel approach that is intended to address this challenge.

An automatic security document classifier system should involve pre-processing, text understanding, comprehending multi-level security grades followed by the application stages. The system that we propose is through an approach that realizes the degree of importance of given text via a rhetorical tree. It then tags the respective sections with the appropriate security level according to the security standards (policies set by organization) that is defined priori. For the purpose of realizing the importance of portions of documents, we make use of the rhetorical structure theory (RST) [2, 3].

The remaining of the paper is organized as follow: Section 2 highlights related work. Section 3 introduces the concept of the proposed technique and gives an overall description of the system. Section 4 discusses an experimental study that was conducted on the system and the results obtained. The experiment was conducted on Arabic texts. Section 5 concludes the paper and highlight future work.

2. Related Work

Little work and limited research has been made in the area of security classification of documents. The most updated research [11] incorporates section level security classification of documents only, a top-bottom approach. Damiani et al. [12] proposed an access control model for XML documents. The model is confined to DTD (Data type definition) level only. In this model, each DTD is associated with particular information that is contained in the document, and this is used to decide which part can and cannot be accessed by user.

Information Security Management Tool, developed by the University of Auckland (New Zealand) for the security protection of the local University official documents, fragments documents into classes with certain implementation of security levels ranging from top to bottom). The NIST, National Institute of standards and Technology (Under Secretary of Commerce for Technology, United States), Tool was developed by NIST for the information protection in documents trafficking between their mother institute and its child institutes. Microsoft Office Document Classifier [13] classifies and labels Office documents, using

document markings that identify the existence of confidential and private information. This tool helps in stopping information leakages of hidden content in Microsoft Word documents, and encourages proper handling of sensitive information. Titus Labs Document Classification (Titus Labs Document Classification for Microsoft Office) [14] is a classification tool for Microsoft Office that allows government and military customers to manage the classification, distribution and retention of valuable corporate documents. Users are constrained to select from a dropdown menu the appropriate classification labels for their documents, spreadsheets, or presentations. ArticSof (Cryptographic tool) [15] is used for email classification, encryption and digital signature. SECLORE (Document's right management tool) [16] uses dynamic rights for distributed document usage control.

For the purpose of realizing the importance of portions of documents, we make use of the rhetorical structure theory (RST). RST was originally developed as part of studies of computer-based text generation [2]. It has been developed to serve as a discourse structure in the computational linguistic field. RST gives the coherence in text [4]. It is intended to describe texts, rather than the process of creating or reading and understanding them. It uses a set of rhetorical relations that associate spans of text in an attempt to identify the importance of the portions. The rhetorical relations can be described functionally in terms of the writer purposes and the writer assumptions about the intended reader. These rhetorical relations hold between two adjacent spans of texts (although there are some exceptions).

The output of applying the rhetorical structure theory to a text is a tree structure that organizes the text based on the rhetorical relations [4]. Each relation connecting two spans of a text may be of two cases. In the first case, the relation connects two spans where the first is identified as a nucleus, representing the semantic of the two spans (i.e., this is more important to the reader than the other span), and the second span is called satellite. In the second case, the two

spans have the same importance to the reader. In this latter case, the relation is called multinuclear relation where both participating spans are considered nucleus. The process of parsing the text and building the rhetorical structure is called the rhetorical analysis. During the process of the rhetorical analysis, the elementary units that participate in building the rhetorical schema are determined, and the rhetorical relations that hold among these units are

also determined to connect the two spans. Determining the potential relations that connects the two spans could be done using several techniques; one of which is determining the rhetorical relations through the cue phrases [8]. Marcu [6] has given cue phrases that can be used in the English language processing. The process of building the rhetorical structure may lead to more than one structure, consequent upon the nature of the natural language text that

stipulates that more than one relation could be assumed to connect two spans. At the end, the emergent structures are most likely closed, but in some cases; they may lead to ambiguity.

We selected RST as the basis for our technique as it allows classifying texts on the basis of the importance of its constituents which serves the purpose of assigning security labels to parts of a give document that is available electronically. There are other text classification techniques that have explored in other applications such as in [17, 18, 19].

3. The Proposed Approach

3.1 The Concept

The rhetorical structure that is built from the rhetorical analysis process is represented as a binary tree that connects the two spans of a text [4]. Each node of the tree has a status (referred to as status information) which has the value of nucleus, satellite, a promotion which represents the most important text unit in this sub-tree, or the type, which represents the relation that connects the two spans [4]. The rhetorical structure tree could be built using the algorithm proposed in [8]. Marcu [5, 6] argues that the promotion of the root of the whole tree gives the reader the most important unit(s) in the text. With this in mind, the classifier uses the status information to tag a certain portion of the document with proper security classification such as secret, confidential, etc. That is, the classifier uses the promotion to determine if the document is about some information that belongs to the specific classification. This information could be easily stated as a secret, and so on. It is possible that the promotion could be only one unit; therefore, the classifier may go some levels down if it wants to cover more units. The level the classifier should go depends on the importance of the information.

The proposed technique that could be used to classify the different parts of a certain document parses each paragraph in the document and builds the rhetorical tree that represents its structure. Then it determines what each paragraph is about by examining the promotion of the root of the tree. It uses the promotion to determine if the importance of the paragraph conforms to the user instructions. In such a case, the classifier labels the paragraph with the required classification. The technique could be adopted to work on pages or sections rather than paragraphs. One merely needs to know which text unit the technique could be applied on. The technique can also be adopted to go deeper to lower levels in the rhetorical tree to determine to which class this part of the text belongs. However, the main idea is to extract the promotion of the text unit. In cases such as in very sensitive documents, the user may desire

to intervene in tagging portions of the document with the proper security class. In such cases, the classifier is able to determine the important parts of the text so that the user supplies the proper tags.

3.2 An Overview of the Classifier System

Figure 1 and 2 depicts details of the proposed system that is based on RST. The working of the system follows the following steps:

RST Processor takes as input the text and builds its corresponding RS Tree. Tree Level Selector traverses the RST, level by level starting from the root. Whenever a sentence (a text unit) is requested by the inference engine, the tree level selector picks the sentence from the current level down to the leaf. As the tree is structured according to the most important sentence (nucleus-based), the probability is greater that a best fit of the classification of the text is at the higher level of the tree.

Inference Engine processes the text unit based on the associate rules that are specific to the document type and decides whether the text can be classified at this stage or it requires the process of more text units. A looping communication with the Tree level Selector is established for this purpose. DB Classifier takes the result of the classification from the Inference Engine and update the database accordingly.

RST Processor gets the relations, builds all the valid RS-trees, and then notifies the RS-tree Selector that the RS-trees are ready.

RuleBase: This is a domain specific knowledge base that reflects the rules of the organization for the security classifications of their sensitive documents. It guided the classifier in tagging the proper security class in the absence of the intervention of the users.

RS-trees Selector selects the most suitable RS-tree for Arabic text summarization.

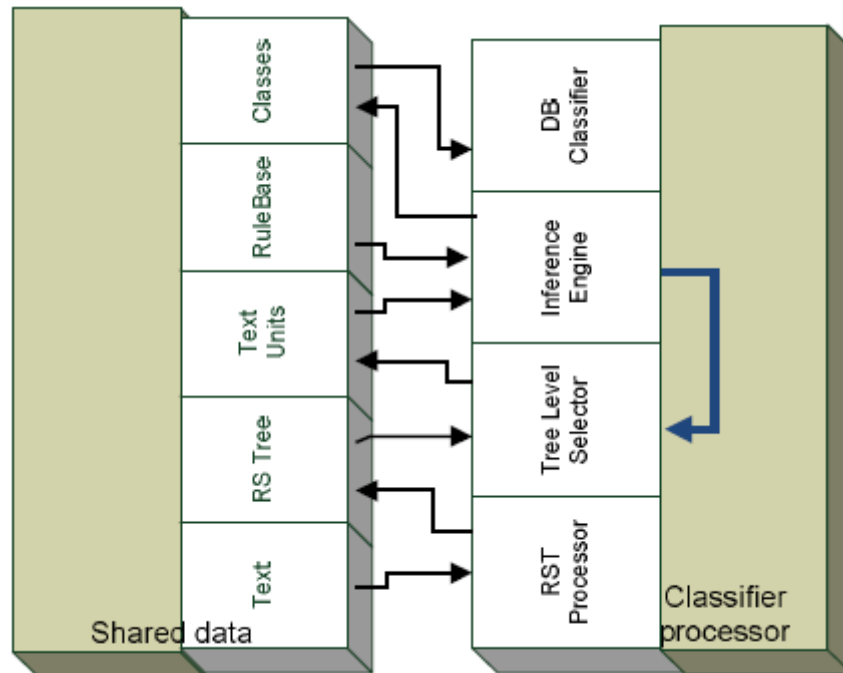


Figure 1 An overview of the classifier system

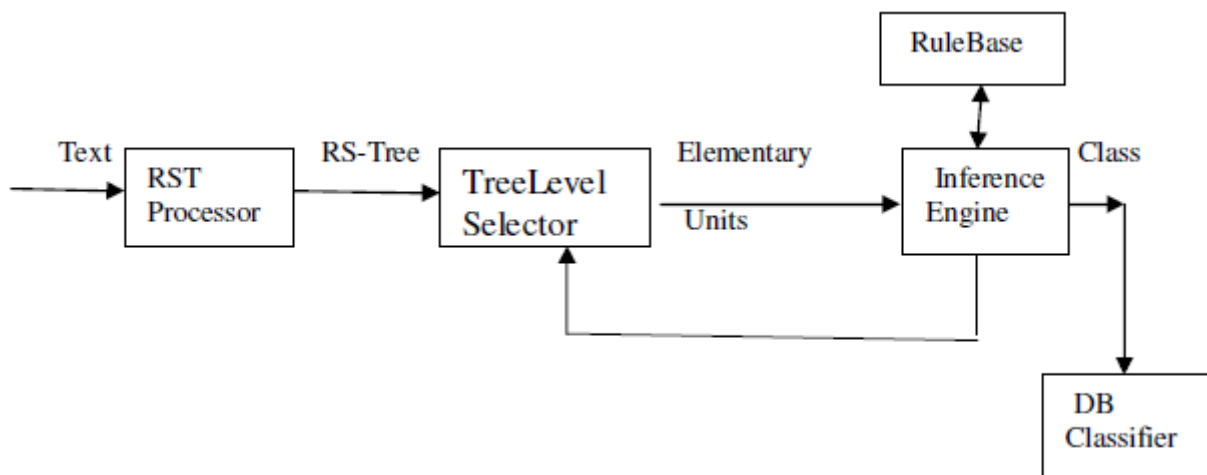


Fig.2 Interactions among the components of the classifier system

3.3 Using the Cue Phrases in the Classifier System

We use the concept of cue phrases [4] to identify the rhetorical relations. The cue phrases were identified for the Arabic text based on the analysis of large Arabic corpus. Part of the RuleBase DB Classifier Inference Engine TreeLevel Selector RST Processor cue phrases and the relations they determine are listed in Table 1. The listed cue phrases are only those that are related to texts used in the experiment for the security classifier of documents. Thus, they constitute only a partial list of the cues that were obtained from the corpus analysis.

Table 1. A Partial list of Cue Phrases

Cue Phrase		Relation
1	و	Joint
2	حيث	Elaboration
3	لذا	Evidence
4	بالرغم	Concession
5	لكن	Contrast
6	على أن	Condition
7	بما أن	Background

Some cue phrases link two spans without respecting the positions these cue phrases appear in. Cue phrases 1 and 5 (Table 1) trigger a multinuclear relation in which both participating spans are nucleus. These two cue phrases always join the span that comes before them and the span that comes after them. The other cue phrases have two cases. They may join the span that comes before them and the span that comes after them, or they may join the span that comes after them (span 1) and the span that comes after span1 –the next two adjacent spans. This, however, depends on the token that comes before the cue phrase. If the token is a dot ("."), a new line ("\\n"), or nothing (i.e. the first sentence in the document), the second case is applied; otherwise the first case is applied. The following example explains the two cases with the cue Phrase (بما أن) (in English since):

[بما أن الشركة وقعت العقد،¹ [فإن العمل في المشروع سيبدأ بعد شهر من الآن]² .

In the above example, the second case of linking the two spans is applied. In this case, the two consecutive spans that come after the cue phrase are linked. Now consider an example of the first case –in which the same cue phrase comes between the two spans it links.

[ستنم دراسة جميع المشاريع المطروحة]¹ [بما أن العمل في المشروع السابق قد أوشك على الانتهاء]² .

In the two examples above, the nucleus and satellite positions differ. In the first example, span ⁽¹⁾ is the satellite and span ⁽²⁾ in the nucleus. In the second example, the nucleus and

satellite exchange their positions –span ⁽¹⁾ is the nucleus and span ⁽²⁾ is the satellite. Therefore, the relations that represent the two examples are:

rhel_rel (Background, 1, 2)

rhel_rel (Background, 2, 1)

The following two tables show the cue phrases in the two cases, and the satellite and nucleus positions in each case.

Table 2. The two spans of the cue phrases appeared in the middle of the paragraph

Cue Phrase	First Span	Second Span
و	Nucleus	Nucleus
حيث	Nucleus	Satellite
لذا	Satellite	Nucleus
بالرغم	Nucleus	Satellite
لكن	Nucleus	Nucleus
على أن	Nucleus	Satellite
بما أن	Nucleus	Satellite

Table 3. Cue phrases that can come at the beginning of the paragraph

Cue Phrase	First Span	Second Span
حيث	Satellite	Nucleus
بالرغم	Satellite	Nucleus
على أن	Satellite	Nucleus
بما أن	Satellite	Nucleus

In a case that a sentence is not joined to any other sentence via a cue phrase, this sentence is joined to the one before it through the relation:

Rhet_rel (Joint, Second_sentence, First_sentence)

For example, in the following text, the parser will consider the second sentence as a case in which the author joins second sentence to the first one.

[التقارير المالية الصادرة من الإدارة المالية تفيد بزيادة المبيعات في هذه السنة.]¹ [هذه الزيادة ناتجة عن الانتعاش

الإقتصادي في السوق المحلية]²

4. Experimental Results

We have performed experiments on Arabic texts to demonstrate the working of classifier. The document units that we applied the classifiers on are the paragraphs. The technique examines each paragraph, determines what it is about, and then classifies it.

The following is an example of how the classifier works. As mentioned above, the classifier works on each paragraph and classifies it. Consider the paragraph below, which describes a company previous contracts and its intent to contract with an expert, and assume that the company considers such information to be classified as secret.

[وقعت الشركة في هذه السنة عدة عقود مع بعض شركات التقنية لتزويدها بالتجهيزات اللازمة].¹ [بالرغم من الجهود المبذولة من قبل إدارة التدريب،]² [إلا أن موظفي الشركة لم يصلوا بعد للمستوى المأمول للتعامل مع هذه التقنية،]³ [لكن الموظفين لديهم دافع للتعليم،]⁴ [إذا ستقوم الشركة بالتعاقد مع خبير تقني ليستفيد منه الموظفون]⁵ [يما أن تكلفة التعاقد أقل من تكلفة الدورات التدريبية،]⁶ [على أن لا تزيد مدة التعاقد على خمس سنوات]⁷ [وتشمل إقامة بعض الندوات العلمية،]⁸ [حيث تثرى هذه الندوات معلومات الموظفين].⁹

According to the description given about the cue phrases used, the following relations hold in the above paragraph:

rhete_rel (Joint, 2, 1)
 rhete_rel (Concession, 2, 3)
 rhete_rel (Contrast, 4, 3)
 rhete_rel (Evidence, 4, 5)
 rhete_rel (Background, 6, 5)
 rhete_rel (Evidence, 7, 6)
 rhete_rel (Joint, 8, 7)
 rhete_rel (Elaboration, 9, 8)

Using the algorithm mentioned in [8], the classifier will build the RST that represents this paragraph. Figure 1 shows a generated tree that the classifier will use to classify this paragraph. The classifier will examine the promotion of the root node to find that sentence (5) is the most important unit in this paragraph. It then uses this fact as a basis to determine its security class and depending on the information obtained from the knowledge base. Since in our example any information regarding the company contracts are to be labeled as secret, this paragraph will be

classified as secret since the promotion indicates that the company is going to make a contract with an expert to train its employees, which is:

لذا ستقوم الشركة بالتعاقد مع خبير تقني ليستفيد منه الموظفين.

If the user configures the classifier to classify this information under another class it will do so; otherwise it will be left unclassified.

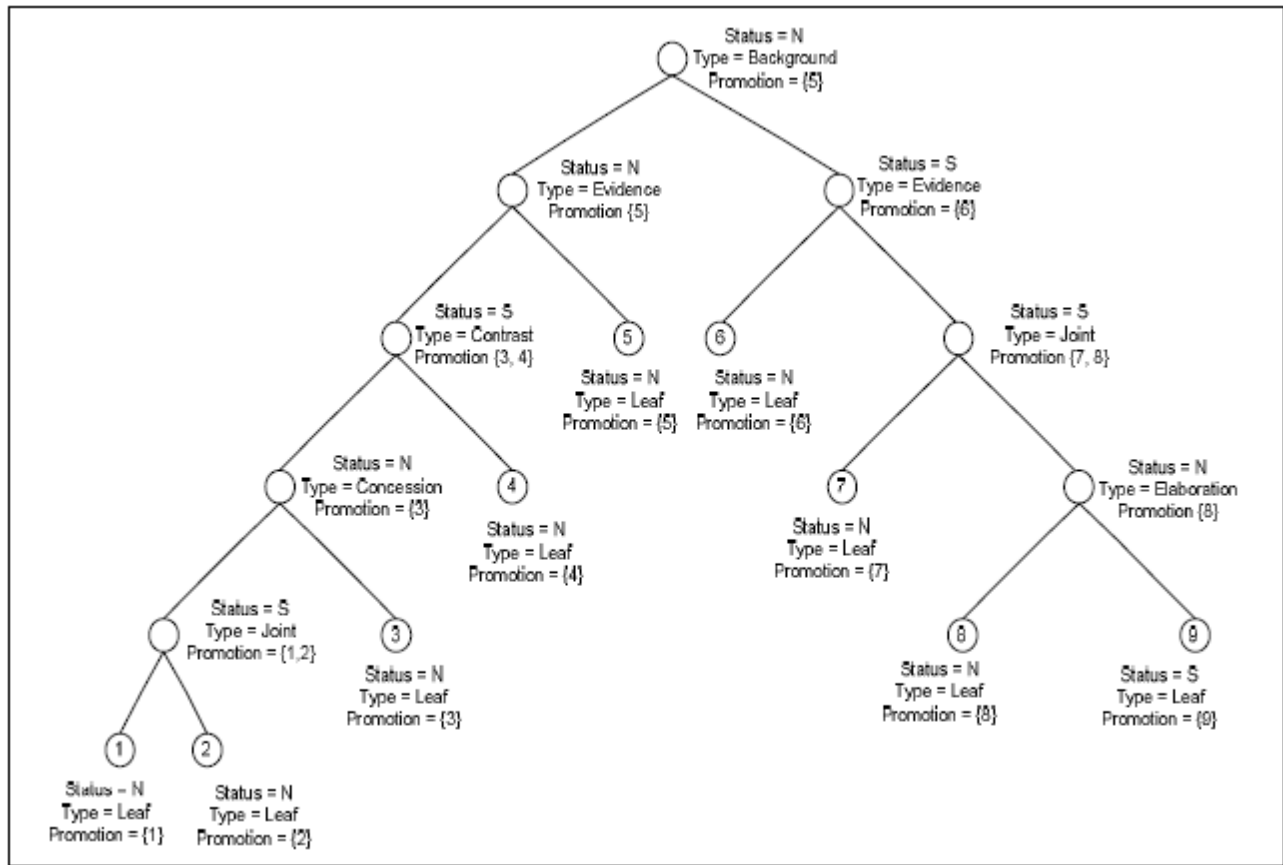


Figure 3. RST representing the test text

5. Conclusion

We have given an introduction to the information classification and its usage. We have shown its importance to ensure the privacy of an organization. The process of classifying the information might be slow and take too much effort when the document is huge. Developing an automated process of classification will definitely lead to faster classification of the documents. The classification process depends on the semantic of the text rather than the syntax; therefore, the technique that should be used to classify the information automatically should use its

semantic. The rhetorical structure theory (RST) is one of the techniques that try to extract the semantic of the text. We have proposed a classification technique that uses this theory to extract the semantic of the text and then classifies the text. We have performed a experiments on certain Arabic texts, and has produced an optimistic result. Currently, we are investigating the units of a document (paragraph, page, section) that could be used to build the rhetorical tree and then classify it. The depth level that the classifier goes for in looking for the promotions is a potential subject of future research. The research should also state the language the technique is being applied to, since the nature of one language may differ from another one. Techniques applied on certain language might not be applicable (or difficult to apply) in another language. We are also investigating the possibility of augmenting the RST techniques with text classification techniques such as statistical and lexical cohesion techniques to attain an improved outcome.

6. Acknowledgement

This work is partially supported by the research center of the college of computer and information sciences in King Saud University.

7. References

- [1] W. Nicolls, Implementing Company Classification Policy with the S/MIME Security Label, The Internet Society, 2002.
- [2] William C. Mann, Sandra A. Thompson, Rhetorical structure theory: Toward a functional theory of text organization", Text, Vol. 8, No.3, pp.243-281, 1988.
- [3] Bill Mann, An introduction to rhetorical structure theory (RST), <http://www.sil.org/~mannb/rst/rintro99.htm>, 1999.
- [4] Daniel Marcu, "Discourse trees are good indicator of importance in text", Advances in Automatic Text Summarization, The MIT Press, 1999, pp 123-136.
- [5] Daniel Marcu, Building Up Rhetorical Structure Trees, The Proceeding of the Flexible Hypertext Workshop of the Eighth ACM International Hypertext Conference, 1997.
- [6] Daniel Marcu, From discourse structure to text summaries, The Proceeding of the ACL'97/EACL'97 Workshop on Intelligence Scalable Text Summarization, Madrid, Spain, 11 July 1997, pp 82-88.
- [7] J. Kajava, Information Security Standards and Global Business, Industrial Technology, 2006. ICIT 2006, IEEE International Conference on 15-17 Dec. 2006 pp. 2091 – 2095.
- [8] Daniel Marcu, The theory and practice of discourse parsing and summarization, The MIT press, UK, 2000.
- [9] M. Gupta, Security Classification for Documents, Computers and Security, Volume 15, Number 1, 1996, pp. 55-71.

- [10] R. W. Maule, Enterprise knowledge security architecture for military experimentation, IEEE International Conference on Systems, Man and Cybernetics, Oct. 2005, Volume 4, pp.3409 – 3414.
- [11] M. Alhammouri, and S. Muftic, A Design of an Access Control Model for Multilevel-Security Documents, 10th International Conference on Advanced Communication Technology (ICACT 2008), Feb. 2008, Volume 2, 17-20.
- [12] Ernesto Damiani, Sabrina de Capitani di Vimercati, Stefano Paraboschi, and Pierangela Samarati, Design and implementation of an access processor for xml documents, The International Journal of Computer and Telecommunications Networking, Vol. 33, June 2000, pp. 59-75.
- [13] <http://www.re-soft.com/product/titus-classification-office.htm>
- [14] http://www.titus-labs.com/software/DocClass_default.html
- [15] <http://www.articsoft.com/>
- [16] <http://www.seclore.com/>
- [17] Rıdvan Saraçoğlu, Kemal Tütüncü, Novruz Allahverdi, A new approach on search for similar documents with multiple categories using fuzzy clustering, Expert Systems with Applications, Volume 34, Issue 4, May 2008, pp. 2545-2554.
- [18] Songbo Tan and Jin Zhang, An empirical study of sentiment analysis for Chinese documents, Expert Systems with Applications, Volume 34, Issue 4, May 2008, Pages 2622-2629.
- [19] Rey-Long Liu, Interactive high-quality text classification, Information Processing & Management, Volume 44, Issue 3, May 2008, pp. 1062-1075.
- [20] W. H. Baker, Is Information Security Under Control?: Investigating Quality in Information Security Management, Security & Privacy, IEEE, Volume 5, Issue 1, Jan.-Feb. 2007, pp.36 – 44.
- [21] Thomas R. Peltier, Information Security Policies, Procedures, and Standards: guidelines for effective information security management, Auerbach publications, Boca Raton, FL 2002.

Visual Inspection of Ceramic Tiles Surfaces Using Statistical Features and LVQ of Artificial Neural Networks.

Ibrahim El-Henawy

Dept. of Computer Science, Faculty
of Computer and Information Sciences,
Zagazig University, Zagazig, Egypt,
Henawy2000@yahoo.com

Samir Elmougy

Dept. of Computer Science, College of
Computer and Information Sciences, King
Saud Univ., Riyadh 11543, Saudi Arabia,
mougy@ccis.ksu.edu.sa

Ahmed El-Azab

Dept. of Computer Science, Misr for
Engineering and Technology(MET)
Academy, Mansoura, Egypt,
ahmed.elazab@yahoo.com

Abstract: A product inspection is an important stage in the product manufacturing process before packing process. Visual inspection is an application of computer vision and image processing. Using Automatic Visual Inspection (AVI) in ceramic tile inspection gives more reliable, fast, robust, and less human intervention, increases processing stability, and improves overall production performance. Many approaches had been applied for this purpose starting from primitive pixel by pixel comparison to advanced techniques such as statistical, structural, and signal processing techniques. In this paper, we use statistical approach as a feature extraction approach for visual inspection of ceramic tiles surfaces defects with neural network classification techniques. Our results show that statistical techniques used for extracting ceramic tile features with LVQ neural networks classifiers give subtle results especially LBP and GLCM which give 96% and 90% of correct classification respectively.

Keywords:- Visual Inspection, neural networks, statistical approach, structural techniques, signal processing techniques, ceramic tiles.

1. Introduction

A product inspection is an important stage in product manufacturing process before packing process. It is applied to many types of products such products like wood [1], steel [2], textile [3], stone [4], plastic products [5] and ceramic tiles [6, 7]. All stages of ceramic tiles industry are done in automatic processes except the final stage of the manufacturing process where it is still performed manually to sort tiles into distinct categories or to reject those found with much defects and faults. This manual process may not give acceptable results because the human inspectors may vary in judgment according to tiredness, or changing in lighting conditions. Boukouvalas justified the commercial and safety benefits for using Automatic Visual Inspection (AVI) in ceramic tile inspection as follows [6]:

- Giving more robust and less costly inspection.
- Producing a significant reduction of human intervention in dusty and unhealthy environment.
- Saving time and efforts in manual inspection.
- Increasing the processing stability and improving overall production performance.

A typical visual inspection system held on production line shown in Figure (1) [7].

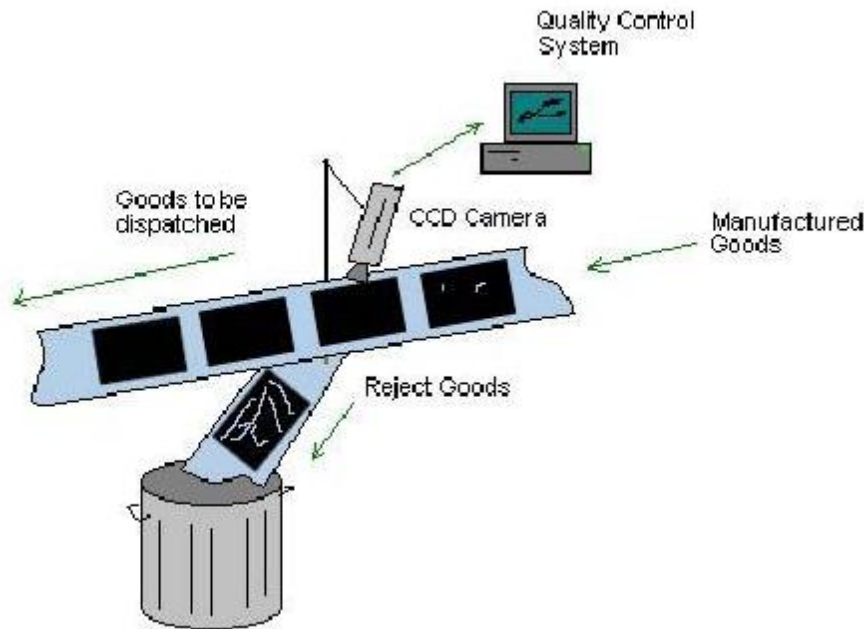


Figure 1. A typical visual inspection system held on production line

In this paper, we consider the statistical feature extraction approach. We use histogram properties, co-occurrence matrix, LBP and LVQ neural network for supervised pattern classification. This paper is organized as follows. Overview of visual inspection process is discussed in Section 2 followed by a short discussion for each step. Proposed feature extraction is presented in Section 3. LVQ neural network structure and its algorithm explained in Section 4. Experimental results are shown in Section 5. Finally, conclusion and future work in Section 6.

2. Visual inspection process

There are four main steps in visual inspection process as shown in Figure (2). These steps are image acquisition, preprocessing, feature extraction, and classification (Recognition). The first step of visual inspection process is image acquisition which means capturing the image of the tile using a CCD (Charge-Coupled Device) camera held on production line. The second step is pre-processing which means making the captured image suitable for processing and feature extraction by applying some operations such as contrast stretching, rotation and thresholding on the captured image. The third and fourth steps are explained in details in Section 2.1 and 2.2 respectively.

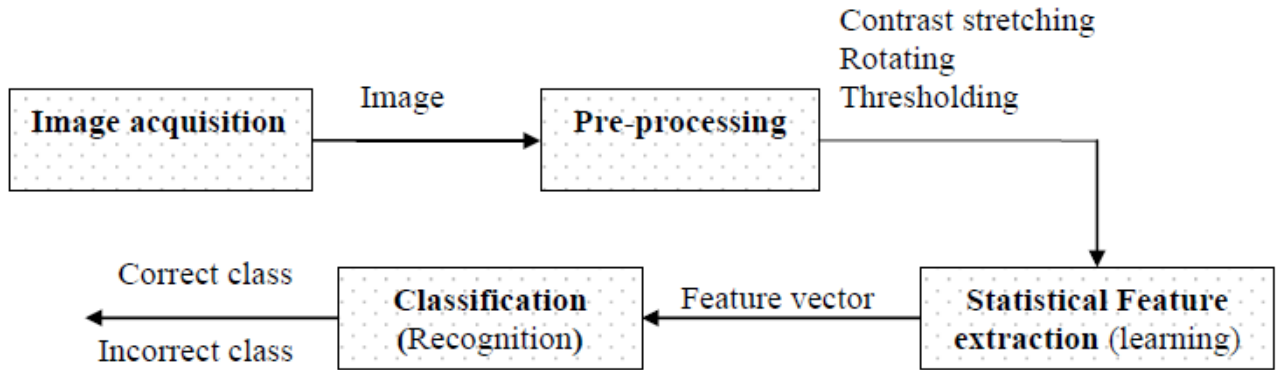


Figure 2. Inspection process steps

2.1 Feature extraction

For pattern recognition, a set of feature is extracted from the image to be classified (learning) and to be assigned to a specific class using classifiers. There are many techniques for feature extraction and textural defect detection. Tuceryan and Jain [8] divided these techniques into four approaches as the following:

(a) Statistical Approach

In this approach, spatial distribution of gray values is analyzed by computing local features at each point in the image, and derives a set of statistics from the distributions of the local features. Depending on the number of pixels, defining the local feature statistical methods can be further classified into a first-order (one pixel), second-order (two pixels) and higher-order (three or more pixels) statistics. In the first-order, statistics estimate properties (e.g. average and variance) of individual pixel values with ignoring the spatial interaction between image pixels. In the second and higher-order statistics, estimate properties of two or more pixel values occur at specific locations relative to each other. The most widely used statistical methods are co-occurrence features and gray level differences [9].

(b) Geometrical methods

In this approach, texture is considered to be composed of texture primitives with attempting to describe the primitives and the rules governing their spatial organization. The primitives may be extracted by edge detection with a Laplacian-of-Gaussian or difference-of-Gaussian filter by adaptive region extraction, or by mathematical morphology [6]. Image edges and contours are often used as primitive elements.

(c) Model-based methods

Model based methods depend on hypothesizing the underlying texture process, and then constructing a parametric generative model which could have created the observed intensity distribution. The intensity function is considered to be a combination of a function to represent the known structural information on the image surface and an additive random noise sequence [10].

(d) Signal processing methods

In this type of techniques, some methods such as Fourier transform Cosine transform and Wavelet transform are used to analyze the frequency content of the image. There is another view for feature extraction is given by Xianghua Xie [11]. He divided the approaches to statistical (histogram properties, co-occurrence matrix, local binary pattern, autocorrelation and registration-based), structural (primitive measurement, edge features, skeleton representation and morphological operations), filter based (spatial domain filtering and frequency domain filtering), and model based approaches (fractal models, random field model and texem model).

2.2 Classification

Classification means assigning a specified pattern to a class amongst different classes according to similarity criteria. These classes may be previously determined or not. If the class previously determined the classification is said to be supervised and unsupervised otherwise. The most known classification techniques are artificial neural network (such as LVQ and Back propagation), and statistical techniques (Maximum likelihood and k-Nearest Neighbors).

3. Feature Extraction using Statistical Approaches

As we discussed before, spatial distribution of gray values is analyzed by computing local features at each point in the image, and derives a set of statistics from the distributions of the local features. In the following subsections, we discuss the statistical methods that we use in this paper.

3.1 Histogram Features of Images

Histogram of a digital image with L total possible intensity level in the range $[0, G]$ is defined as the discrete function:

$$H(r_k) = n_k \quad (1)$$

where r_k is the k -th intensity level in the interval $[0, G]$, n_k is the number of pixels in the image whose intensity level is r_k and, G is the number of gray scale colors equal 255. The normalized histogram is obtained by dividing all elements of $h(r_k)$ by the total number of pixels in the image [12] which is denoted by n as:

$$p(r_k) = \frac{h(r_k)}{n} = \frac{n_k}{n} \quad (2)$$

3.2 Co-Occurrence Matrix

Another important method in statistical approach is the Gray Level Co- Occurrence Matrix (GLCM) which is one of the earliest texture feature extraction method and widely used nowadays introduced by Haralick [13, 14]. GLCM is a second order statistics rather than histogram, which is first order statistics. It is an $N \times N$ matrix represents the relation between the pixel and other pixels adjacent to it in a specific direction or displacement and it is defined as in Equation [3]:

$$Cd(i,j) = |\{(r,c): I(r,c)=I \text{ and } I(r+dr, c+dc)=j\}| \quad (3)$$

Although there is a problem with driving texture measure from the co-occurrence matrices when choosing the displacement vector d . Zucker [15] showed that the solutions for this problem is to use a χ^2 statistical test given in Equation [4] to select the value(s) of d that have the most structure; that is to maximize the value:

$$X2(d) = \sum_i \sum_j \frac{Nd2[I_i, I_j]}{N_d[i] N_d[j]} \quad (4)$$

After computing the gray level co-occurrence matrix, normalizing co-occurrence matrix is required by using the following [Equation 5]:

$$N_d[i,j] = \frac{C_d[i,j]}{\sum_i \sum_j C_d[i,j]} \quad (\text{between 0 and 1}) \quad (5)$$

Normalizing the co-occurrence values to lie between zero and one and allow them to be thought as probabilities in a large matrix is given in Equation [5]. As soon as $N_d[i, j]$ is computed, the important features can be extracted from it.

3.3 Local Binary Pattern (LBP)

LBP was first introduced by Ojala et al in 1996 [16] as a shift invariant complementary measure for local image contrast. LBP operator works with the 8-neighbours of the pixel and using the value of the center pixel as a threshold. An LBP code for a 3×3 neighborhood is produced by multiplying the thresholded values with weights given to the corresponding pixels, and summing up the result and computing the contrast of the image by subtracting the average of the gray levels below the center pixel from the average of gray levels above (or equal to) the center pixel. As shown in Figure [4].

LBP has the average of a simple computation and an invariant to image rotation [11]. Also LBP can be computed according to various neighborhood.

Sizes and distances 4×4 or possibly 5×5. In basic form of this method, LBP operator in one neighborhood of the image is defined as in Equation [6]:

$$LBP_{P,R} = \sum_{i=0}^{P-1} s(g_i - g_c) 2^i, \quad s(x) = \begin{cases} 1, & x \geq 0, \\ 0, & x < 0. \end{cases} \quad (6)$$

Where the gray value of the central pixel is g_c and the gray value of the i th neighbor is g_i . According to this definition, it is seen that the output of the operator is P bit binary number with 2^P distinct values.

4. Feature Extraction of Ceramic Tiles using Statistical Approaches

In this paper, we apply the statistical feature extraction on 100 gray scale images of resolution 256×256. Features extracted from these images are then compared using LVQ neural networks. The key reason of using statistical feature extraction is the speed of computation and robustness of capturing properties [2, 9, 11, 16]. In the following subsections, we discuss our using and

implementing of histogram, Gray Level Co-Occurrence Matrix, and Local binary pattern as methods of feature extraction for ceramic tiles.

4.1 Histogram Method

Figures [3a] and [3b] show the image of both normal and defective tile with the defects shown in [3c]. The corresponding histograms for Figures [3a] and [3b] are shown respectively in Figures [3c] and [3d]. The features extracted from the histogram are mean (Equation [7]), variance (Equation [8]), skewness (Equation [9]), kurtosis (Equation [10]) and many others. In this paper, we use histogram features shown in Table [1].

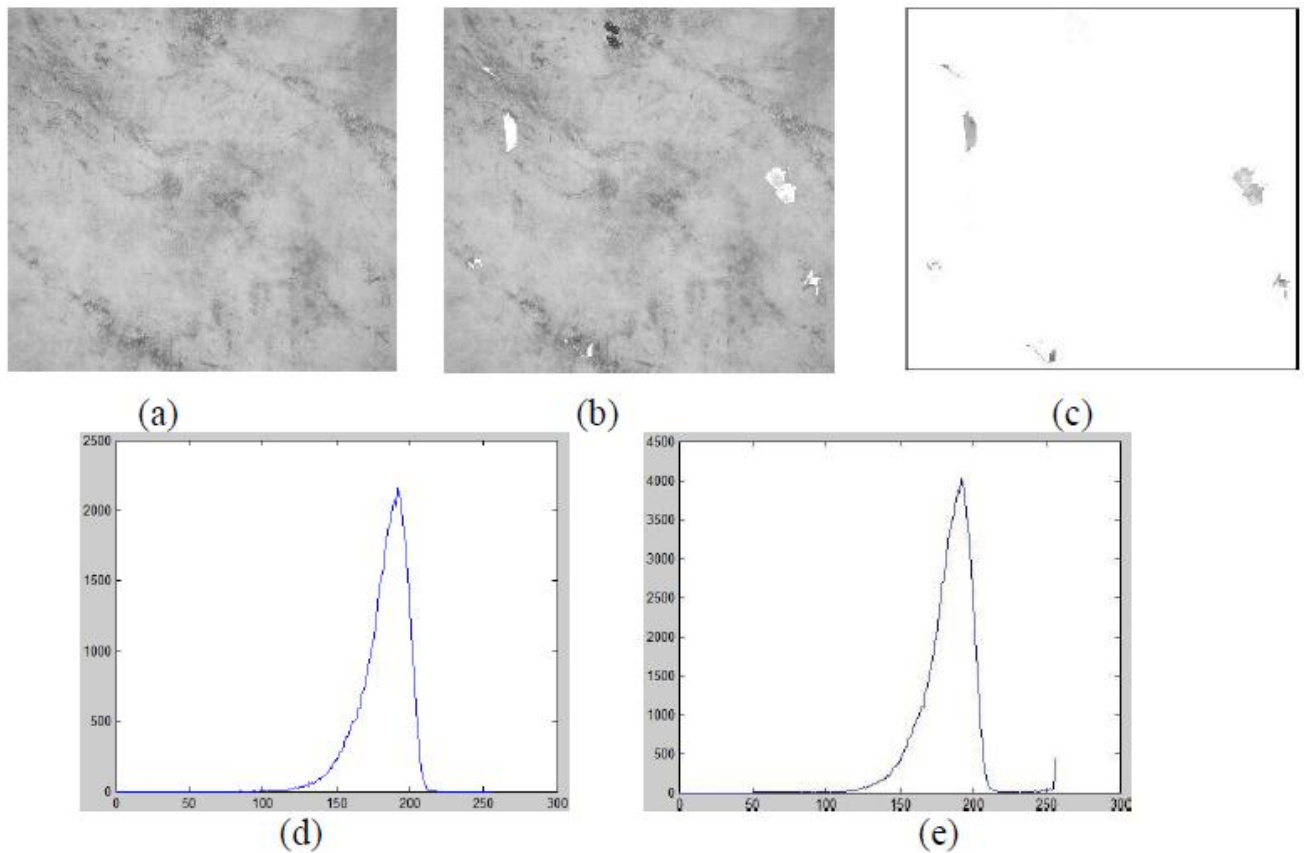


Figure 3. A normal tile (a), defective tile (b), the defects of image b (c) and corresponding histogram for each one ((d) and (e))

Table (1) Histogram, LBP features.

Mean	$\bar{x} = \frac{1}{n} \cdot \sum_{i=1}^n x_i$	(7)
Variance	$\sigma^2 = \frac{\sum (X - \mu)^2}{N}$	(8)
Skewness	$skew(X) = \frac{E[(X - \mu)^3]}{\sigma^3}$	(9)
Kurtosis	$kurt(X) = \frac{E[(X - \mu)^4]}{\sigma^4}$	(10)

4.2 Co-Occurrence Method

The Co-Occurrence matrix computed from the ceramic image is used to extract some features. First feature extracted from GLCM is entropy (Equation [11]) that measures the disorder of an image texture. The second feature is energy (Equation [12]) that measures uniformity of an image texture. The third feature is homogeneity (Equation [13]) that measures the local homogeneity of a pixel pair. The fourth feature is contrast (Equation [14]) that measures local contrast of an image and Table [2] shows the features of GLCM where a vector of these features is used for the classification phase.

Table [2] GLCM Feature

Feature name	Equation
<i>Entropy</i>	$Entropy = -\sum_i \sum_j N_d[i, j] \log_2 N_d[i, j]$ (11)
<i>Energy</i>	$Energy = \sum_i \sum_j N_d^2[i, j]$ (12)
<i>Homogeneity</i>	$Homogeneity = \sum_i \sum_j \frac{N_d[i, j]}{1 + i - j }$ (13)
<i>Contrast</i> (<i>Inertia</i>)	$Inertia = \sum_i \sum_j (i - j)^2 N_d[i, j]$ (14)

GLCM on 4×4 of a sample of sub-image we are using is shown in Figure [4] and the corresponding features are shown in Table [3].

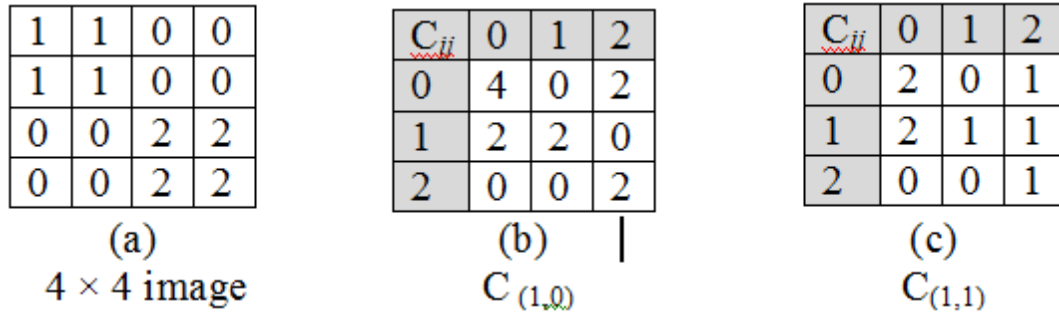


Figure 4. 4 × 4 sub image (a), GLCM at C (1, 0) and C (1, 1) (b) and (c)

Table [3] GLCM features for sub-image in Figure (4)

	$C_{(1,0)}$	$C_{(1,1)}$
Energy	0.22222	0.18750
Entropy	-2.25121	-2.50103
Inertia	0.83333	0.87513
Homogeneity	0.80555	0.97294

4.3. Local Binary Pattern Method

LBP and Contrast (C) are computed for the ceramic tile as shown in Figure [4]. The implementation in Figures [5a, 5b] show LBP for normal and defective tiles of Figure [3a, 3b] while Figures [5c, 5d] show Contrast for normal and defective tiles of Figures [3a, 3b].

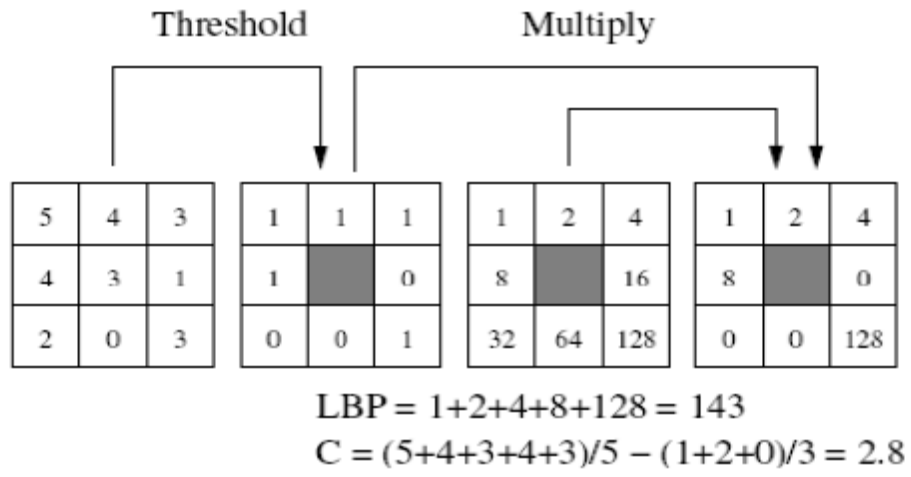


Figure 4. Computing the LBP and Contrast from [17]

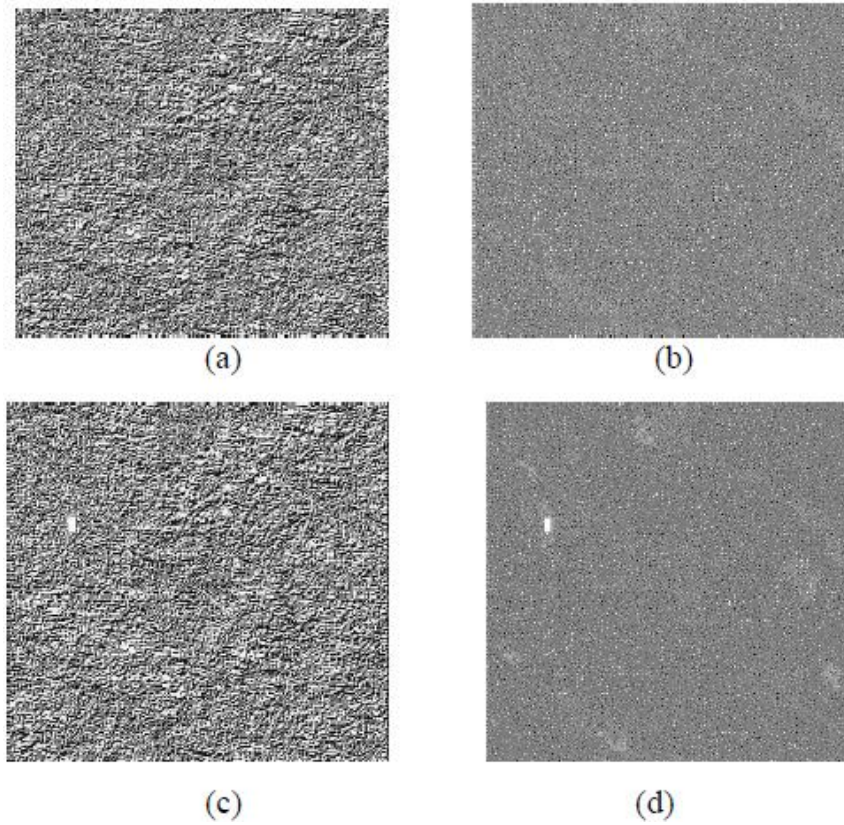


Figure 5. LBP and Contrast for tiles in Figure (3)

5. LVQ structure

We use LVQ neural network for ceramic tiles classification. LVQ is a supervised version of the Self-Organizing Map (SOM) algorithm by Kohonen in 1986 [18]. The strength of LVQ over other supervised neural network techniques, like Back Propagation, is trained significantly faster and can handle data with missing values. The LVQ neural network structure shown in Figure [6] consists of two layers. The first layer called the input layer which contains the neurons for the input vectors while the second layer called output layer consists of two neurons representing the two classes; the first one for normal (defect free) tile and the other for abnormal (defected) ones.

The main idea of LVQ is that for each input vector, calculate the neuron that is closest to it using a distance measurement, like Euclidian distance Equation [15], to determine the winner neuron and then update its weights by "winner-take-all" principle [19]. LVQ neural network algorithm is shown in Figure [7].

$$d(i,j)=\sqrt{(|x_{i1}-x_{j1}|^2+|x_{i2}-x_{j2}|^2+...+|x_{ip}-x_{jp}|^2)} \quad (15)$$

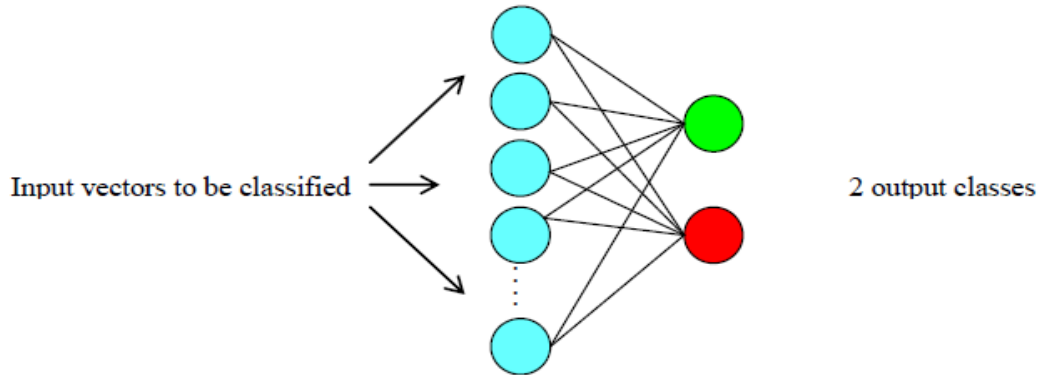


Figure (6) LVQ structure

Step 0: Initialize the reference vectors.

Step 1: Repeat until no. of training reached **OR** when all classes are true

Step 1.1: For each training input pattern **x**

-Step 1.1.1: Find Euclidian distance that is minimum among all output nodes

-Step 1.1.2: Update **wj** based on class matches the input vector's class **x**.

* If $\text{class}(\mathbf{x}) == \text{class}(\mathbf{w}_j)$ then $\mathbf{w}_j = \mathbf{w}_j + \alpha(\mathbf{x} - \mathbf{w}_j)$.

* If $\text{class}(\mathbf{x}) \neq \text{class}(\mathbf{w}_j)$ then $\mathbf{w}_j = \mathbf{w}_j - \alpha(\mathbf{x} - \mathbf{w}_j)$.

Step 1.2: Reduce the learning rate

Figure (7): LVQ algorithm.

6. Experimental Results

LVQ network was trained three different times on our tiles data set images using different features extracted from the previous techniques. The first time training by using the histogram extracted features which listed in Table (3) as the following:

$$\text{Histf} = \{\text{mean, variance, skewness, kurtosis}\}$$

Second time training by using GLCM extracted features as the following; Coocf = {energy, contrast, entropy, homogeneity} Third time training by combining both of LBP/ contrast extracted feature as the following:

$$\text{LBP/Contrastf} = \{\text{LBP_mean, LBP_variance, LBP_skewness, LBP_kurtosis, Contrast_mean, Contrast_variance, Contrast_skewness, Contrast_kurtosis}\}$$

Classification Accuracy (CA) is computed using the following equation:

$$CA = \frac{\text{No.of} \cdot \text{correct.tiles}}{\text{Total} \cdot \text{No.of} \cdot \text{tiles}} \times 100 \quad (16)$$

Features	CA (%)
Histogram	84
GLCM	90
LBP/Contrast	96

6. Conclusion and Future work

Statistical techniques used for extracting ceramic tile features with LVQ neural networks classifiers give subtle results, especially GLCM and LBP that gives 90% and 96% of correct classification respectively. These results show that LBP is the better statistical feature extraction method than GLCM and Histogram methods. One of the proposed future works is to use statistical and signal processing techniques for inspecting color ceramic tiles with computation speed up.

7. References

- [1] F. Pernkopf. "Detection of surface defects on raw steel blocks using Bayesian network classifiers". Pattern Analysis and Applications, 7:333–342, 2004.

- [2] O. Silvén, M. Niskanen, and H. Kauppinen, “Wood inspection with non-supervised clustering”, *Machine Vision and Applications*, 13:275–285, 2003.
- [3] I. Rossi, M. Bicego, and V. Murino. “Statistical classification of raw textile defects”, In *IEEE International Conference on Pattern Recognition*, volume 4, pages 311–314, 2004.
- [4] J. Liu, and J. MacGregor, “Estimation and monitoring of product aesthetics: Application to manufacturing of engineered stone countertops”. *Machine Vision and Applications*, 16(6):374–383, 2006.
- [5] T. Petkovic, J. Krapac, S. Loncaric and M. Sercer, "Automated visual inspection of plastic products", *Proceedings of the 11th International Electrotechnical and Computer Science Conference ERK 2002*, pp. 283-286, Portoroz, Slovenia, 2002.
- [6] C. Boukouvalas, J. Kittler, R. Marik, M. Mirmehdi, and M. Petrou, “Ceramic tile inspection for colour and structural defects”, *Proceedings of AMPT95*, ISBN 1 872327 01 X, pp. 390–399, August 1995.
- [7] H. M. Elbehery, A. A. Hefnawy, and M. T. Elewa, “Visual Inspection for Fired Ceramic Tile's Surface Defects Using Wavelet Analysis”, *GVIP(05)*, No. V2, pp. 1-8, January 2005.
- [8] M. Tuceryan, and A. Jain, “Texture Analysis”, *Handbook of Pattern Recognition and Computer Vision*, chapter 2, pages 235–276, World Scientific, 1998.
- [9] M. M. Trivedi “Object Detection Based on Gray Level Co-occurrence”. *Computer vision, Graphics, and image processing*, vol. 28, pp: 199-219, 1984.
- [10] S. W. Zucker and D. Terzopoulos, “Finding Structure in Co-occurrence Matrices for Texture Analysis”. *Computer graphics and image processing*, vol. 12; pp: 286- 308, 1980.
- [11] X. Xie. “A Review of Recent Advances in Surface Defect Detection using Texture analysis Techniques”. *Electronic Letters on Computer Vision and Image Analysis* 7(3):1-22, 2008.
- [12] R. C. Gonzalez, R. E. Woods, and S. L. Eddins, "Digital Image Processing Using Matlab", PEARSON Prentice Hall, 3rd Ed, pp 76-88, 2004.
- [13] R. Haralick, K. Shanmugam, and I. Dinstein, “Texture features for image classification”, *IEEE Transactions on Systems, Man, and Cybernetics*, SMC- (6), pp: 610-621, 1973.
- [14] A. Monadjemi, “Towards Efficient Texture Classification and Abnormality Detection”, PhD Thesis, University of Bristol, UK, 2004.
- [15] S. W. Zucker and D. Terzopoulos, “ Finding Structure in Co-occurrence Matrices for Texture Analysis” *Computer graphics and image processing*, vol. 12; pp: 286-308, 1980.
- [16] T. Ojala, M. Pietikainen, and T. Maenpaa, "Multiresolution Gray-Scale and Rotation Invariant Texture Classification with Local Binary Patterns", *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 24(7):971–987, 2002.
- [17] T. Maenpaa, and M. Pietikainen, "Texture analysis with local binary patterns", *Handbook of Pattern Recognition and Computer Vision*, pages 197–216, World Scientific, 3rd Edition, 2005.
- [18] D. T. Pham, and S. Sagioglu, "Neural network classification of defects in veneer boards", *Professional Engineering*, Vol 214, p.255-258, 2000.

- [19] Laurene Fauseett, “Fundamentals of Neural Networks Architectures: Algorithms and Applications”, Florida Institute of Technology. Prentice hall, pp: 187-195. 1994.

Evaluation of MUVES: Needs and Results

Samir M. Abd El-razek

Dept. of Information Systems, Faculty
of Computer Science and Information
Systems, Mansoura University, EGYPT

Hazem M. EL-bakry

Dept. of Information Systems,
Faculty of Computer Science and
Information Systems, Mansoura
University, EGYPT

Waeil Fathi Abd El-wahd

Head of Operations Research and
Decision Support Systems, Faculty of
Computers and Information,
Menoufiya University, Egypt

Abstract: *The use of virtual patients in medical education is increasing rapidly. The integration of the e-learning modules is essential for their success. Quality Assurance measures have been applied to MUVES; it has been assessed using some different testing types. Also a collection of virtual patients of MUVES were introduced and introduction sessions were organized, and then divided students into small groups. Two studies were done using questionnaires and log file analysis, both studies measured student's satisfaction, learning outcome, design of virtual patients and the perception of different integration scenarios.*

Results:

Key words: *Virtual Environment, Virtual Patient, Case Based Learning, Medical E-Learning.*

4. Introduction

The medical colleges in Mansoura and Monofeya Universities have a main focus of interest in introduction of new teaching methodologies. One of these new teaching methodologies is Virtual Patients (VPs). VPs are interactive computer systems that simulate real life clinical scenarios [1]. VPs allow students to get standardized experiences and knowledge with clinical problems. Students can be introduced to uncommon diseases regardless of real patient availability. Virtual patient systems are interactive learning environments in which students can learn relevant skills without risk for real patients, like clinical decision making. Training can be repeated with these systems as needed and feedback should be provided to improve skills of students.

Evaluation in e-learning systems used to gather information about the impact or effectiveness of this Web-based learning event [2]. Measurements might be used to improve the offering, determine if the system objectives were achieved, or determine if the offering has been of value to the organization [3].

This paper presents the evaluation results of MUVES, a virtual environment for medical e-learning in medicine college of Mansoura. MUVES is using the problem based approach, in which students are choosing what is the next step in the process of clinical case diagnosis. They have to gather information from history questions, physical examination, laboratory and technical tests, and then they should be able to make appropriate decision accordingly. MUVES provides a user-friendly interface and necessary tools to generate virtual patients and play virtual patients.

Section 2 of this paper gives a short overview about MUVES and its functionality, then the rest of paper describes the evaluation result of MUVES. Section 3 describes MUVES quality

evaluation for functional system testing results and some nonfunctional measures like performance. Section 3 presents teachers evaluation and then section 4 presents students evaluation, which are both based on a survey conducted in the two medical colleges. Respondents of this survey have given feedback on the current and future use of MUVES, including different educational settings and scenarios within which virtual patient have been used. Finally section 5 presents the usability analysis of the system based on system logs analysis, which presented many important indicators of medical elearning systems design and implementation.

5. Overview of MUVES

MUVES is a medical e-learning system for authoring and delivering CBL to medicine students based on MVP standards [4]. MUVES can work as helper tool for medicine teachers to demonstrate the diagnosis process for students over the web in a virtual lecture style. Teachers can select or create a new VP template, these templates are categorized by teacher upon their relevance to subjects. Each VP template contains a complete VP data and media resources, the teacher can create an instance from this template to start performing the diagnosis process on this instance. The idea of creating templates is to “reuse” previously created VPs and make it easier for creating VPs rather than copying files. After creating the instance the teacher may need to change some VP data for more clarification on this case or to add new important information needed for diagnosis like up-normal values for some lab elements, or increase the measure of blood pressure, etc. An example of the authoring module is shown in figure 1.

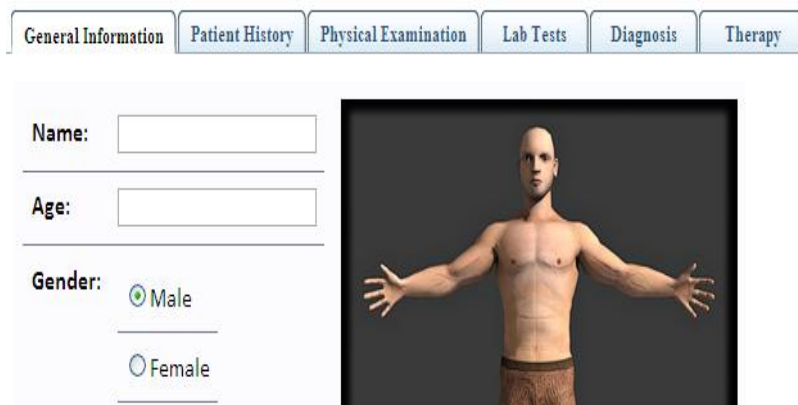


Figure 1 VP authoring Module

All the previous steps are performed “offline” or away of students. Once the teacher finishes his updates on the instance chosen, he can start the virtual lecture. The virtual lecture items are broadcasted to students over the web, starting with VP data and media resources. Teacher starts his lecture using collaboration tools like audio/video streaming, whiteboard, and performing

visible physical examination on the VP instance. The physical examination is performed using a rich set of physician tools, and possibly to acquire some lab results for this VP instance. Each step of the diagnosis process can be opened for discussion by students using a simple chatting tool with the teacher until the teacher reach the conclusion of the diagnosis student can submit their questions and discuss the reasonability of diagnosis. The lecture can be saved to the Virtual Lecturers Library (VLL) and can be reviewed by student later. Another important feature of MUVES is it can be used as an assessment tool for students. Teacher can “setup” the instance as introduced previously, but he will assign the diagnosis process to a group of students who need to collaborate together to conclude the possible diagnosis with guidance of the teacher.

6. MUVES Quality Evaluation

MUVES has been applied to a full test cycle before publishing. The testing aimed to enhance the quality of the system by insuring that all functional and nonfunctional requirements have been met. Figure 2 shows the progress of closing defects found in system testing along different versions of the system

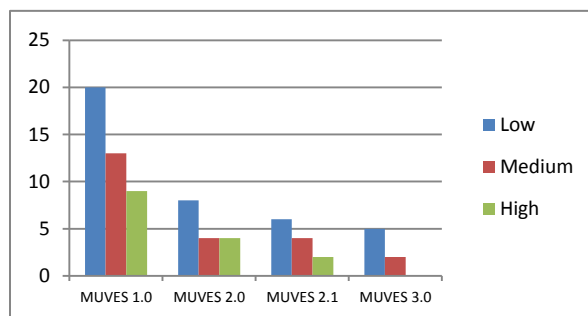


Figure 2. MUVES Versions Bugs Report

Good progress has been achieved in decreasing the number of bugs, especially ones of high severity which represents bugs of type unplanned error screen or fatal errors. MUVES 3.0 which is the current version still has some minor errors like some JavaScript errors which are caused by using development tools like JQuery and AJAX. The continuous bug findings are caused by the continuous additions and enhancements to the system.

Another type of testing was performed to assess some functional requirements like performance testing. The primary goal of performance testing is to study application performance under normal load and heavy load for sustained durations or for short durations. The reports and graphs generated at the end of the test session should provide important indicators for application stability under heavy load, and response time of the system is reasonable.

An important part of a load test process is to analyze the errors and adjust the tests in order to achieve meaningful and accurate results from the test. The following reports and graphs give detailed indicators for system performance.

Time-based frequency of errors over the duration of the test. Error rate shown in figure 2 is an important metric in stress testing. This indicates the maximum number of users that can be served correctly, without errors and hence your site capacity. You will also need to watch error rate during the test to verify that the error is within acceptable range even after a long run.

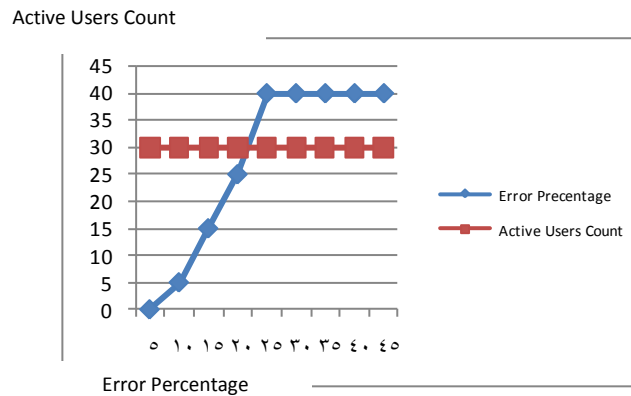


Figure 3. Error Rate Graph: Time Vs Error

Graph shown in figure 4 indicates average response time for the each page through time. Each bar in the graph is the average server response time for each page. From this graph it possible to identify peaks in the response time for critical pages. Ideal behavior of response time graph is that response time does not increase with load. The point at which the graph increases sharply indicates beyond this load server cannot serve the request and users will see no response or a very slow response.

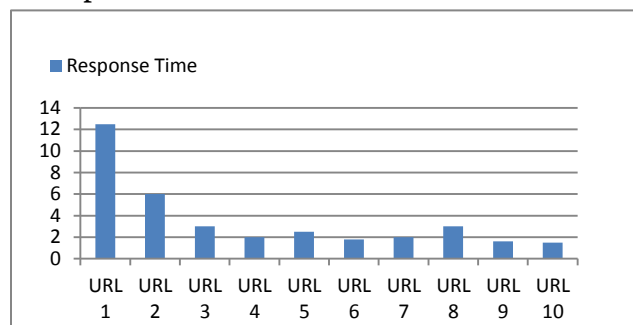


Figure 4. Response Time Graph

7. Teachers' evaluation

This evaluation was conducted by interview with a group of teachers to capture their detailed feedback on the process of creating VPs and integration with standard systems in a qualitative manner. Teachers commented that their creating was done in two steps.

Firstly, teachers would create the existing cases from international references, there is a range of differing healthcare cultures from most rural to most urban). In addition to this the teacher would also check the validity of units, reference ranges for laboratory tests, and healthcare NHS protocols like the National Institute of Clinical Excellence guidelines. This was a straight forward process and took approximately an hour per case as there weren't many changes in how a patient is treated and cared for between local and foreign.

Secondly, the VPs needed to be story-boarded and expanded. This was done during an initial brainstorm meeting between two subject matter experts.

A few questions have been asked to teachers to evaluate the idea and the system.

2.1 In general, is there a place for VPs in the curriculum? Please explain your answer.

Every form of teaching targets a specific area. The traditional fact-based approaches don't allow the students to learn in a realistic environment, to make mistakes and by doing so to become more skilled in the clinical decision making process. The contact with the real patients on the wards is very limited and the prevalence of some important diseases is too small to give the opportunity for each student to come face to face with it and learn how to proceed and how to make the right decisions. That's why there is a gap for realistic and interactive virtual patients in the curricula which are always accessible to the students. Many people learn by doing, for some it's also much more easier to learn in this way. Students can't learn how to make decision with other resources. Students have to be able to make mistakes.

2.2 How do VPs compare with other types of resources that teachers create for the student? Do they differ in quality from other types of learning resources used by teachers?

The VPs and their quality is so different that they're not really comparable to other learning resources. They're targeting other parts of curriculum. To speak about the quality of a VP, we need to define how this quality should be estimated.

2.3 On average, how long did it take you to create a case?

First stage (brainstorming): about 1hr, Second stage: writing the content and creating the pathways was very time consuming, Total time needed: 9 hours.

2.4 What experience is needed in order to create VPs effectively?

Native language speaker and experienced professional clinician. It helps to have been in the situation to make the VP realistic (To be able to find enough alternative paths)

2.5 What do you think are the best features of?

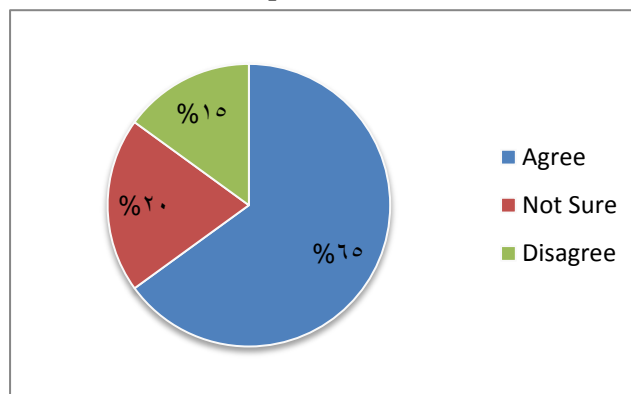
You really feel like you're in the case – but more in a tutorial way.

8. Students survey results

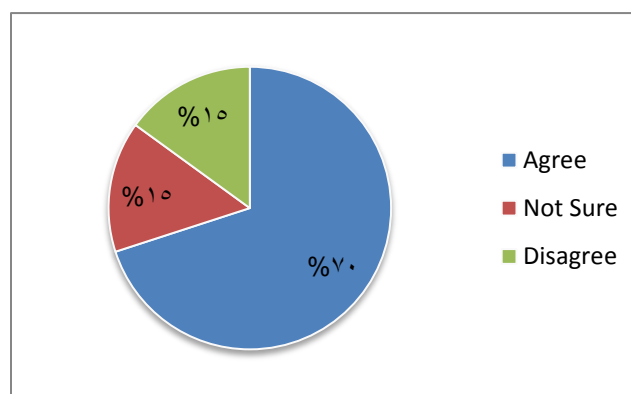
A group of students from medical faculties were selected to be trained on using MUVES. After conducting training sessions, students worked with three medical cases. Virtual patients were developed by teachers, required support was provided to help them on the authoring tools. All students were divided into small groups, each group was using one computer. Sessions durations were about 2 hours, and were led by teachers and research team. After working through the VPs, students were asked to fill in questionnaire comparing virtual patients as a new method of teaching with the classic methods of teaching. The classic methods consisted of bed side teaching and theoretical lectures.

This is a questionnaire completed by 12 students who completed individual questionnaires following the virtual patients that were delivered weekly in their module. Below are the results:

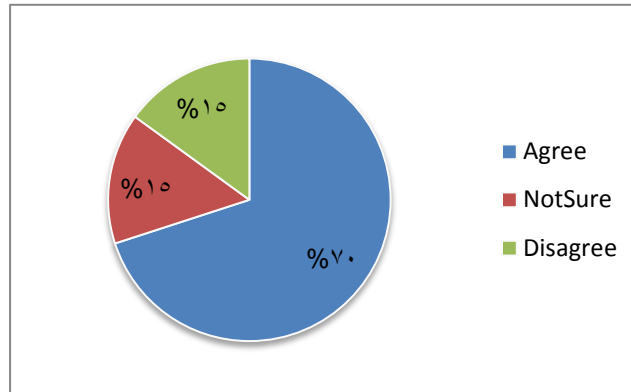
A- After completing this case, I feel better prepared to confirm a diagnosis and exclude diagnosis in a real life patient with this complaint.



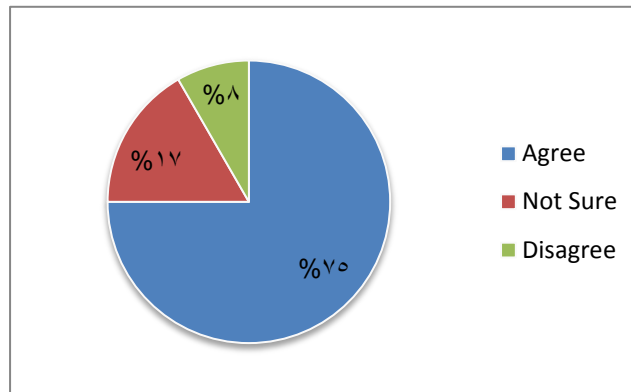
B- The feedback I received was helpful in enhancing my diagnostic reasoning in this case



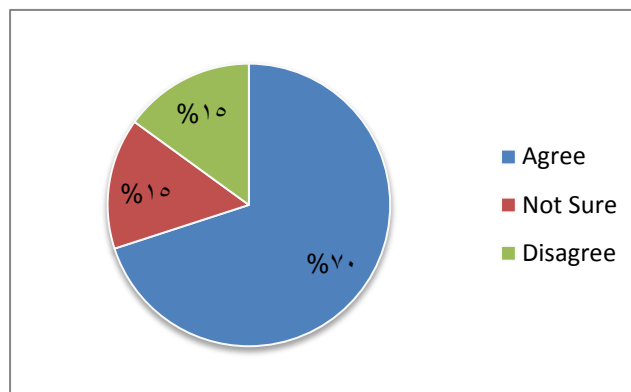
C-After completing this case I feel better prepared to care for a real life patient with this complaint.



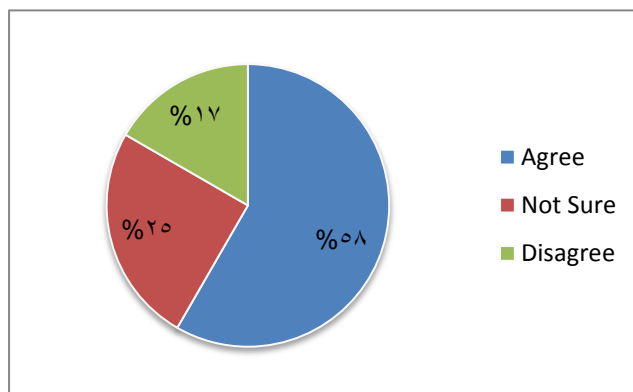
D- The questions I was asked while working through this case were helpful in enhancing my diagnostic reasoning in this case.



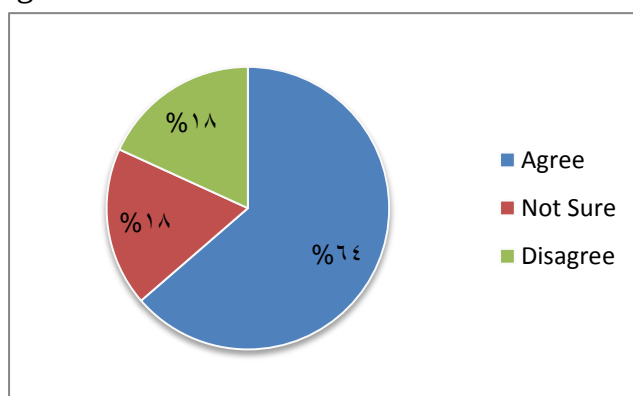
E-While working through this case, I was actively engaged in revising my initial image of the patient's problem as new information became available.



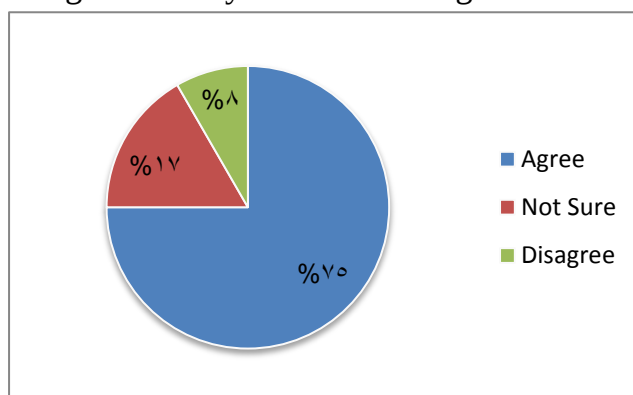
F-While working on this case, I felt I had to make the same decisions a doctor would make in real life.



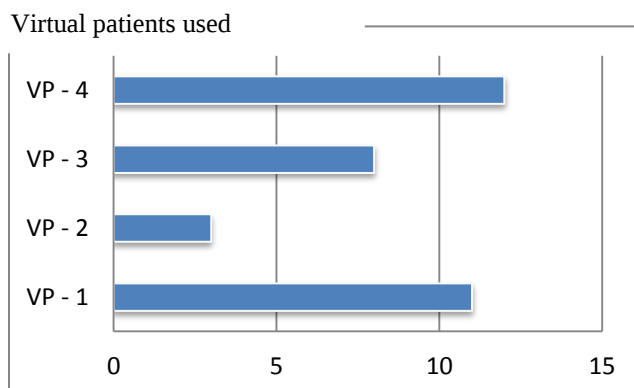
G-While working through this case, I was actively engaged in creating a short summary of the patient's problem using medical terms.



H-While working through this case, I was actively engaged in thinking about which findings supported or refuted each diagnosis in my differential diagnosis.



I-Overall, working through this case was a worthwhile learning experience.



9. Usability Analysis

During the whole time all activities of the participants in MUVES were recorded in the form of a Log file. The participants produced between 11273 and 26158 lines in the log file. The navigational behavior of each participant was analyzed. MUVES has five main navigation forms: a tree structure (similar to the Windows Explorer), breadcrumbs (a description of the hierarchical path on top of the main text field), a sitemap and links within the text. In addition to the main forms of navigation, additional factors have been taken into consideration: number of clicks on links leading to a virtual lecture, number of clicks on virtual exams. The last two categories might also be seen as text links but we think their function differs considerably from normal text links, therefore they have been analyzed separately. Table 1 contains a minimum and a maximum number for all categories. These numbers show the maximum (minimum) number of clicks in this category for a specific user, the average number of clicks in this category and the percentage across all users. The values in the min- and max-lines do not necessarily come from one single user.

These data also show that the tree navigation was most popular among users. It is probably so attractive because all users are well acquainted with this form of navigation and it can be used easily because it is always visible on the screen (in contrast to the sitemap which has to be accessed specifically). The other methods of navigation were only used rarely.

Table 1: Usage of Various Navigational Practices. Navigation Bar (NAV), Tree (TR), Breadcrumbs (BC), Link within Text (TXT), Sitemap (SM), Virtual Lecture (VL), Virtual Exam (VE).

	NAV	TR	BC	TXT	SM	VL	VE
Min.	84	437	0	72	1	63	12
Max.	289	2262	482	883	479	191	112
Avg.	186.5	1349.5	241	477.5	240	127	62
Result	7%	50%	9%	18%	9%	5%	2%

The average number of clicks for breadcrumbs (241.0) is due to one single user who used this form of navigation 482 times. The other users only adopted this form of navigation between 0 and 11 times. The same user also used the sitemap comparatively often. He clicked 479 times on the sitemap to access single pages.

10. Results

Surveys' results have been summarized into the following subjects so that they can inform current and future developments based on student feedback:

11. Motivation

The feedback from the teachers and students suggests that they found the use of virtual patients (VPs) highly motivating. They found the interactive nature of the VPs particularly motivating. This was highlighted by the following quotes from one of the students when comparing the use of VPs with traditional textbook learning:

"It just makes learning a bit more interesting, rather than just learning from a textbook. Because it's always more helpful going in and seeing patients and seeing cases, and this is a way to do that when there isn't a patient, or when it's in the evening and you just want to sit in bed with your Laptop. I have to say I've got books at home but because it doesn't take you through step by step, it just asks you general questions, I find it a lot duller than this. I'd rather use this."

Another general consensus from the focus group was that the students felt that VPs were quick and convenient to use and learn from. They added that it was something that they were easily able to integrate into their studying routine at home. This in turn increased their motivation for using VPs. This was highlighted by the following quote from one of the students:

"I think whenever I want to do some revision or study it always seems like a big thing, I need to set aside at least an hour, whereas with this, I just did it. I had ten minutes and I was a bit bored. I was online, I wanted to go on Facebook, and then there was nothing left to do on

Facebook; I had gone through all my friends. I don't know what else to do. And so I clicked on the VP and I went through it. So it's just being able to slot it in whenever you want. Because I spend a lot of time on the Internet and so it was something I could do when I was bored."

The students also found the use of pictures within the VPs very motivating because it made the learning experience more personal. Especially when treating VPs that they were able to put a face to the name. In fact, the example below shows the impact that killing a VP can have on a student. The inclusion of pictures in VPs is something that the students felt they would remember more than just reading about it in a textbook. This was highlighted by the following quotes from one of the students:

"I think that having the picture of the baby makes it a bit more personal. Oh, it's a real person I'm dealing with rather than just a case in a book. You'll be sat there, or I'll be sat there at an exam with a specific question, and then you think, I can remember this because I remember seeing a case of it. And that's why I think this could potentially be such a useful tool"

12. Current use of VPs

The feedback from the students suggests that they used the VPs in a variety of different ways. Some of the students used the VPs to practice their clinical decision making in a safe environment. This was highlighted by the following quotes from one of the students:

"I like having the opportunity to make clinical decisions and practice my clinical reasoning. Sometimes when I wasn't sure what the answer was, I'd just have a guess and then see what it said. So just sort of learning by what it told me about what was right and wrong."

Students also used the VPs as a tool for self-assessment, highlighted by the following quote from one of the students:

"I used it in the same way as you, in that I would do it without a book to test my own knowledge on the signs, symptoms, presentation, and management of the case."

The students used the VPs as an alternative to a traditional textbook. Some students used VPs as an alternative to learning the same facts that they would acquire from a textbook. Others felt that the VPs brought the facts they learnt in textbooks alive. This is highlighted by the following quotes from students:

"I used it as an alternative to my textbooks. So if I wanted to do some work, I find it more interesting than sitting just with bland textbook facts."

"I didn't sit there with a textbook next to it; I just did it to see what I could remember without the textbook."

"I don't think books always give you a very good sense of what's actually important in clinical practice. Because it's just all dry facts in books, and you can get a bit bogged in the minutia and not really realize, when a patient's in front of you, which bits are relevant at the time and which ones aren't. So I think it helps with that."

Students also found the VPs useful when preparing for their examinations, both to identify gaps in their knowledge and for general revision purposes. This is highlighted by the following quotes from students:

“I think with me – and this is just the way I would work – it’d be something I’d be doing towards whenever I had an exam or a test.”

“It’s also just good to remind you what you need to revise and what you already know, and so you can prioritize your learning.”

13. Future use of VPs

The feedback from the students suggests that they would indeed use virtual patients (VPs) in the future. They felt that they could use the VPs in different ways (i.e. alongside lectures and in combination with textbooks) and at different stages of their learning (i.e. before exams as a revision tool or as a learning tool over the course of a module). This was highlighted by the following quote from one of the students:

“I think I’d read some, do the theory, and then use it. I think I’d probably start to base my revision more on that kind of thing, now, whereas what I’ve always done before is just gone over my notes and made little revision notes. I think I’d now be much keener to move to this sort of resource over revising and just making more endless notes.”

They felt that this method of learning was clinically relevant and real. This was highlighted by the following quote from one of the students:

“They’re all cases that, I think, feature in our learning objectives anyway, so they’re all relevant to stuff that we’re going to be examined on. So it just makes a nice break and a slightly different way of working, and just to put it into practice, as well, the clinical decision-making.”

The students felt that VPs were the most useful e-learning resources that they had ever used. This was highlighted by the following quote from one of the students:

“This is the only thing that I’ve found useful, ever. Well, semi-useful. So I think it might inspire me to use online information more, but then I’m just not used to bothering to go online to learn, so whether I’d remember to, I’m not sure. I’d probably need those email reminders or else I’d forget.”

The students felt that the introduction of VPs had changed their way of learning and their study habits. In the past they felt that they were used to a certain way of learning and stuck to that routine. However, since the implementation of VPs, they had changed their way of learning. This was highlighted by the following quote from one of the students:

“I found it quite easy to get into. Because I was quite stuck in my ways from how I did my career. I was like, no, no, this is how I study and that’s it. I’m still quite set in how I write my notes, but as a means of practicing I’ve actually found it has changed how I’ve worked. And I think it’s something that I would definitely stick with.”

14. Suggested improvements

The students suggested two main improvements to the VPs. Firstly, they wanted to combine small investigative tasks such as simple blood tests into one group, rather than having to choose them individually. This is highlighted by the following quote from one of the students:

“Even down to the simple detail of when you order bloods”

Secondly, students wanted to have the option of doing more than one thing at a time. They felt that this would be more clinically realistic rather than exploring the consequence of each option individually, which they found frustrating. For example when given the option of multiple examinations, they would often want to choose to do more than one at the same time. This is highlighted by the following quotes from students:

“That thing of being given a list of four things and you’d do all of them, and I wish I could have just clicked all of them and to get it over there rather than having to go through it step by step.”

“When it would give you a choice of four things, lots of the time you would do all of those four things. It’s not like you would just choose one thing. So if you’re going to do all of those things anyway, then that might be more realistic, in terms of the ability to tick several things all at once.”

Finally, students wanted there to be a greater number and a wider range of VPs to be integrated into the curriculum. This is highlighted by the following quote from a student:

“I know it’s probably in the early stages of this, but although it did link into my learning, I wish we had a lot more cases, because it only covered a small amount.”

15. Conclusion

Once the virtual patients were repurposed and enriched (i.e. the main outputs of the project were complete) a number of different evaluation studies were conducted with different stakeholders. The objectives of the evaluation were concerned less with the processes by which the project went about its activities, and more with gathering information on the ease of adaptation and effectiveness of VPs. The evaluation was primarily to establish the worth of what has been achieved. Given the aims of the project, it was important to capture the experiences of the students, academic staff and subject matter experts to provide data which would provide insight to inform future developments. A mixed-methods approach was employed to provide a cost-effective approach to collecting and analyzing data.

In practice the use of a virtual patient from one healthcare culture to another was an efficient use of time and resources. This study and other studies demonstrated that even though there is often a strong requirement for virtual patients each time a learning resource is used, it can still be worth the time and effort, if the learning resource has sufficient value in the learning process. This learning was the opportunity for decision-making, exploring consequences of

actions, and for safe practice. Students were able to use these resources in a variety of different ways and learning styles, and recognized the value of a resource that improved practice. It clearly personalized their learning. Teachers and students described the outcome as highly successful.

Students believed VPs provided excellent learning, in a context which improves the making of their profession. It provided them with opportunities to practice clinical reasoning, then take decisions and explore the consequences of their decisions. VPs provided them with opportunities for learning by collaborative discussion and in tutorials, but it also provided opportunities for individual safe practice, personal revision and self-assessment.

The students described VPs as quick and easy to use, easily integrated into their study time, and available anytime, anyplace. Without ever using the phrase, students were describing the virtual patient as an excellent tool for personalized learning. There were a few improvements suggested to aspects of the VP structure, notably a request to batch the tests up in the way that clinicians would normally practice; students did not like trawling through series of test one after another. Most improvement requests were easily addressed. A student comment which was fundamental to the purpose of the MUVES project was that there were too few VPs and students wanted more.

16. References

- [1] OLIVIER CURÉ, Evaluation methodology for a medical e-education patient-oriented information. Informa, 2010
- [2] Pohl, Margit; Rester, Markus: Ecodesign: Development and Testing of an E-learning System. kassel university press, 2011.
- [3] Weiming Wang, Tianyong Hao and Wenying Liu. Automatic Question Generation for Learning Evaluation in Medicine, Springer 2010
- [4] Morikawa T, Moriguchi H, Nose T, Ohoka H. Construction and evaluation of e-learning system for medical treatment safety measures. US National Library of Medicine 2009
- [5] Morikawa T, Moriguchi H, Nose T, Ohoka H. The Evaluation of Game-based E-learning for Medical Education: a Preliminary Survey 2010
- [6] The MedBiquitous Consortium: Enabling Collaboration for Healthcare Education <http://www.medbiq.org>. Accessed 22 June 2011. MedBiquitous 2009
- [7] Fallon C, Brown S. E-learning Standards: A Guide to Purchasing and Deploying Standards-Conformant E-learning. Boca Raton: St Lucie Press, 2003.
- [8] B. Ferry, J. Hedberg, & Harper, B. How do preservice teachers use concept maps to organize their curriculum content knowledge? Journal of Interactive Learning Research, 1998.
- [9] S. Hackbarth, The educational technology handbook: a comprehensive guide: process and products for learning. New Jersey: Educational Technology Publications, 1996.
- [10] <http://www.virtualpatients.eu/about/about-evip> , Last Accessed on July 20, 2010.
- [11] Samir M. abd El-razek. Waeil Fathi Abd El-wahd, Hazem El-bakri,

- [12] "MUVES: A Virtual Enviroment System or Medical Case Based Learning" , International Journal of Computer Science and Network security, VOL. 10 No. 9 , September 2010 , pp 159-163

AN EFFICIENT TECHNIQUE FOR SQL INJECTION DETECTION AND PREVENTION

Esraa Mohamed Safwat
Computer science
Department
Faculty of computer and
information
Menoufyia University
Esraa_safwat87@yahoo.com

Hany Mahgoub
Computer science
Department
Faculty of computer and
information
Menoufyia University
H_mghguob@yahoo.com

ASHRAF EL-SISI,
Computer science
Department
Faculty of computer and
information
Menoufyia University
ashrafelsisim@yahoo.com

Arabi Keshk
Computer science
Department
Faculty of computer and
information
Menoufyia University
arabikeshk@yahoo.com

Abstract: With the recent rapid increase of interactive web applications that employ back-end database services, a SQL injection attack has become one of the most serious security threats. This type of attack can compromise confidentiality and integrity of information and database. Actually, an attacker intrudes to the web application database and consequently, access to data. For preventing this type of attack different techniques have been proposed by researchers but they are not enough because most of implemented techniques cannot stop all type of attacks. In this paper our proposed technique are detection of SQL injection and prevention based on first order, second order and blind SQL injection attacks online. The proposed technique implemented in JAVA and evaluated for seven types of SQL injection attacks. Experimental results have shown that the proposed technique is efficient related to execution time overhead. Our technique need to be one second overhead to execution time. Moreover, we have compared the proposed technique with the popular web application vulnerabilities scanner techniques. The most advantages of proposed technique Its easiness to adopt by software developer, having the same syntactic structure as current popular record set retrieval methods.

Keywords: - Web application, Database SQL Injection Attack, detection, prevention

1. Introduction

Web sites are dynamic, static, and most of the time a combination of both. Web sites need protection in their database to assure security. An SQL injection attacks interactive web applications that provide database services. These applications take user inputs and use them to create a SQL query at run time. In an SQL injection attack, an attacker might insert a malicious SQL query as input to perform an unauthorized database operation. Using SQL injection attacks, an attacker can retrieve or modify confidential and sensitive information from the database.

The popular solutions for the prevention of SQL injection attacks include coding best practices, input filtering, escaping user input, usage of parameterized queries, implementation of least privilege, white list input validation [2,3,4], These solutions should be employed usually during the development of an application. This is the major limitation of such solutions as they do not cover the millions of Web applications already deployed with this vulnerability. Detection of SQL injection and prevention technique is proposed based first order, second order and blind SQL injection attacks online. SQL injection detection and prevention (SQLI-

DP) technique is implemented in JAVA and evaluated for five types of SQL injection attacks. Experimental results have that proposed technique is efficient related to execution time overhead, it's need to be approximately 1second overhead to execution time. In addition, it is easily adopted by software developer, having the same syntactic structure as current popular record set retrieval methods.

The rest of this paper is organized as follows; section 2 present the problem formulation .section 3 presents the background about SQL injection attacks and the concept of parse tree and presents related work. Proposed SQLI-DP technique and experimental results is presented in section 4. Section 5 conclusion and future work is presented. Section 6 provides the References.

2. Problem Formulation

SQL injection is a technique maliciously used to obtain unrestricted access to databases by inserting maliciously crafted strings to SQL queries via a web application. It allows an attacker to spoof his identity, expose and tamper with existing data in databases, and control databases server with the privileges of its administrator. There is a variable SQL injection scanner but it prolongs the connection time if it wants detect or prevent or both. The popular solutions for the prevention of SQL injection attacks can't solve the problem of legacy system and the software developer not easy to adapt. In this paper we try to improve the execution time and try to found way to improve the legacy system to able to prevent injection attacks. Try to make our technique easy to adapt by software developer and make it more efficient in the future.

3. Background of SQL injection attack and parse tree concept

3.1 classifications of SQL injection attacks (SQLIA):

There are different methods of attacks which depend on the goal of attacker are performed together or sequentially [5, 6, 7]. For a successful SQLIA the attacker should append a syntactically correct command to the original SQL query. Show the types of SQLIAs attacks with examples as the following.

Table 1. Types of SQL injection attacks

Types of Attack	Description
Tautologies	SQL injection codes are injected into one or more conditional statements so that they are always evaluated to be true. Exp: Generated SQL Query: <i>SELECT username, password FROM clients WHERE username = 'user1 OR '1'='1 —' AND password = 'whatever'.</i>
Logically Incorrect Queries	Using error messages rejected by the database to find useful data facilitating injection of the backend database. Exp: <i>"Microsoft OLEDB provider for SQL Server (0x80040E07)Error converting nvarchar value „CreditCards” to a column of data type int"</i>
Union Query	Injected query is joined with a safe query using the keyword UNION in order to get information related to other tables from the application. Exp: <i>SELECT * FROM Table1 WHERE id = -1 UNION ALL SELECT null, null, NULL, NULL, convert(image,1), null, null,NULL, NULL, NULL, NULL, NULL, NULL, NULL, NULL, NULL, NULL--</i>
Stored Procedure	Many databases have built-in stored procedures. The attacker executes these built in functions using malicious SQL Injection codes. Exp: <i>SELECT accounts FROM users WHERE login= 'doe' AND pass=' '; SHUTDOWN; -- AND pin =</i>
Piggy-Backed Queries	Additional malicious queries are inserted into an original injected query. Exp: <i>SELECT * FROM products WHERE id = 10; DROP TABLE members;--</i>
Alternate Encoding	the injected text is changed in order to evade detection by defensive coding practices and most of the automatic prevention techniques. Encodings such as hexadecimal, ASCII and Unicode character encoding can be used for attack strings. Exp: <i>SELECT * FROM Accounts WHERE user='user1'; exec(char(0x736875746466f776e)) -- ' AND pass=' ' AND eid=</i>
Blind Injection	An attacker derives logical conclusions from the answer to a true/false question concerning the database. - Information is collected by inferring from the replies of the page after questioning the server true/false questions. Exp: <i>index.php?id=1 and 1=(SELECT 1 FROM information_schema.tables WHERE TABLE_SCHEMA="blind_sqli" AND table_name REGEXP '^[\a-n]' LIMIT 0,1)</i>

3.2 Parse tree concept

Parse tree is a data structure for the parsed representation of a statement. Parsing a statement requires the grammar of the statement's language. By parsing two statements and comparing their parse trees, we can determine if the two queries are equal. When a malicious user successfully injects SQL into a database query, the parse tree of the intended SQL query and the resulting SQL query do not match. By intended SQL query, we mean that when a

programmer writes code to query the database, she has a formulation of the structure of the query. The programmer-supplied portion is the hard-coded portion of the parse tree, and the user-supplied portion is represented as empty leaf nodes in the parse tree. These nodes represent empty literals. What she intends is for the user to assign values to these leaf nodes. A leaf node can only represent one node in the resulting query, it must be the value of a literal, and it must be in the position where the holder was located. By restricting our validation to user-supplied portions of the parse tree, we do not hinder the programmer from expressing her intended query. An example of her intended query is given in Figure 1. This parse tree corresponds to the example we presented in Section 1, `SELECT * FROM users WHERE username=? AND password=?`. The question marks are place holders for the leaf nodes she requires the user to provide.³ While many programs tend to be several hundred or thousand lines of code, SQL statements are often quite small. This affords the opportunity to parse a query without adding significant overhead.

Suppose that [8] a database contains name and password fields in the users table, and a web application contains the following code to authenticate a user's log in.

```
sql = "SELECT * FROM users WHERE name = '" + request.getParameter(name) + "' AND
password = '" + request.getParameter(password) + "'";
```

This code generates a query to obtain the authentication data from database. If an attacker inputs `" or 1=1 -- 1"` into the name field, the query becomes:

```
SELECT * FROM users WHERE name = '' or 1=1 -- 1 AND password = 'xxx';
```

The WHERE clause of this query is always evaluated to be true, and thus an attacker can bypass the authentication, regardless of the data inputted in the password field. Web applications commonly use SQL queries with client-supplied input in the WHERE clause to retrieve data from a database. By adding additional conditions to the SQL statement and evaluating the web application's output, you can determine whether or not the application is vulnerable to SQL injection. For instance, many companies allow Internet access to archives of their press releases. A URL for accessing the company's fifth press release might look like this:

<http://www.thecompany.com/pressRelease.jsp?pressReleaseID=5>

The SQL statement the web application would use to retrieve the press release might look like this (client-supplied input is underlined):

```
SELECT title, description, releaseDate, body FROM pressReleases WHERE
pressReleaseID = 5
```

The database server responds by returning the data for the fifth press release. The web application will then format the press release data into an HTML page and send the response to the client. To determine if the application is vulnerable to SQL injection, try injecting an extra true condition into the WHERE clause. For example, if you request this URL . . .

http://www.thecompany.com/pressRelease.jsp?pressReleaseID=5 AND 1=1

. . . and if the database server executes the following query . . .

SELECT title, description, releaseDate, body FROM pressReleases WHERE pressReleaseID = 5 AND 1=1

. . . and if this query also returns the same press release, then the application is susceptible to SQL injection. Part of the user's input is interpreted as SQL code.

A secure application would reject this request because it would treat the user's input as a value, and the value "5 AND 1=1" would cause a type mismatch error. The server would not display a press release.

The process of generation of queries in a dynamic web application can be represented as a function of user's inputs. In this context, SQL injection is any situation in which the user's input is inducing an unexpected change in the output generated by the function. We define two parameters

SQL Statement = SQL(Argi) $(i=1 \text{ to } n)$

Argi $\leftarrow$ Input from user

SQL() $\leftarrow$ function represented by web application

SQL Statement Safe = SQL(Arg Safe i) $(i=1 \text{ to } n)$

Arg Safe i $\leftarrow$ "qqq" or any single token

We require that the application will not allow the user to enter any part of SQL query directly. We define that two statements are semantically equivalent, if they perform similar activities, once they are executed on the database server. So, if we determine that both SQL Statement and SQL Statement Safe are semantically equivalent, then by definition the SQL Statement is bound to have an expected behaviour and there is no possibility for a SQL Injection. Here semantic action implies a particular activity like comparison, retrieval etc., and not the lexical equality. We use this semantic comparison to detect SQL Injection. The semantic comparison is done by parsing each of the statements and comparing the syntax tree structure. If the syntax trees of both the queries are equivalent, then the queries are inducing equivalent semantic actions on the database server, since the semantic actions are determined by the structure of the SQL statement. For example,

Let,

Arg = { α , β }

Arg Safe = {"qqq", "qqq"}

Now,

SQL Statement = SQL(Arg)

= SELECT * FROM 'User_Table'

WHERE user_name = ' α ' AND password = ' β '

$SQL\ Statement\ Safe = SQL(Arg\ Safe)$
 $= SELECT * FROM User_Table$
 $WHERE user_name = 'qqq' AND password = 'qqq'$

The SQL Statement Safe is parsed to produce a syntax tree as shown in figure 1. We can consider two cases of user inputs, first case without any injection and second with injection.

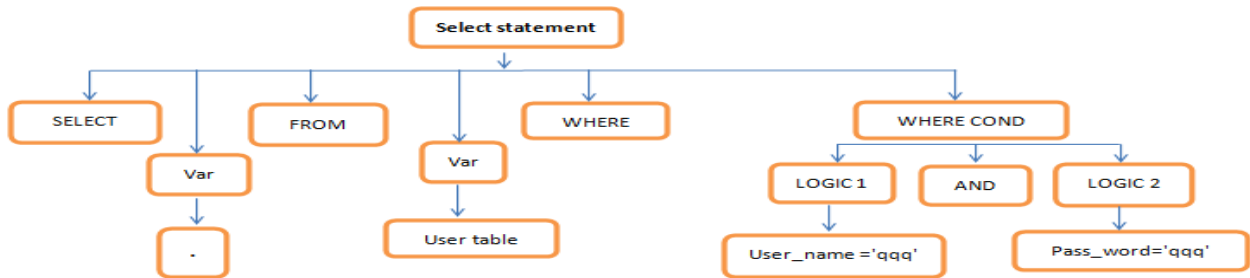


Figure.1. SQL_Statement_Safe

Case 1: Let,
 $Arg\ Normal = \{\alpha, \beta\} = \{admin, admin\ pwd\}$

Now,

$SQL\ Statement\ Normal = SQL(Arg\ Normal)$
 $= SELECT * FROM User_Table WHERE$
 $user\ name = 'admin' AND$
 $password = 'admin_pwd'$

The SQL Statement Normal can be parsed as shown in figure 2. On comparing the semantic structure of SQL statement safe and SQL statement normal, we can see that both of them have similar semantics. Both statements are extracting all values from a table after checking for two logic equalities combined with an AND operator. This implies that there is no possible SQL injection and hence we can safely execute the query.

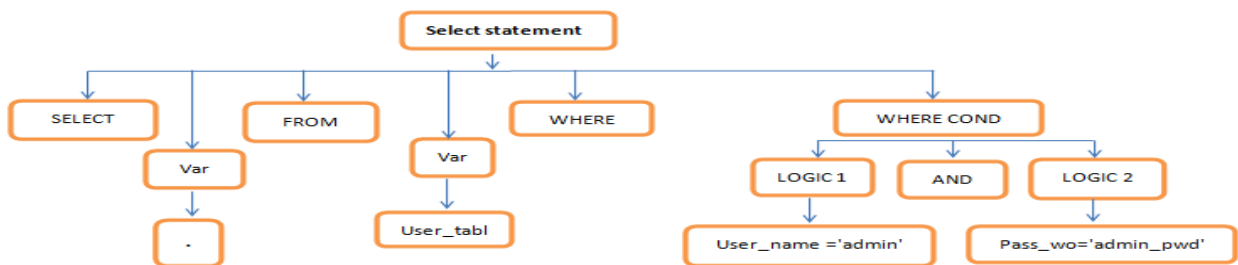


Figure.2. SQL_Statement_Normal

Case 2: Let,
 $Arg\ Injection = \{\alpha, \beta\} = \{admin_OR_1_ =_ 1, hacker\ pwd\}$.

Now,

SQL Statement Injection = *SQL(Arg Injection)*

= *SELECT * FROM Use_Table WHERE user_name = admin OR '1' = '1'*
AND password = 'hacker_pwd'

In this case, on comparison of *SQL Statement Injection* which is represented in figure 3 with *SQL Statement Safe*, we can see that the semantic structures of The two statements are not similar. This is because *SQL Statement Injection* is doing the additional action of "OR '1' = '1' ". This implies that on application of the input, the semantics of the output has been modified. This detects a possible SQL injection and the execution of the query should be stopped [10]. So to prevent and detect the SQL injection we use parse tree.

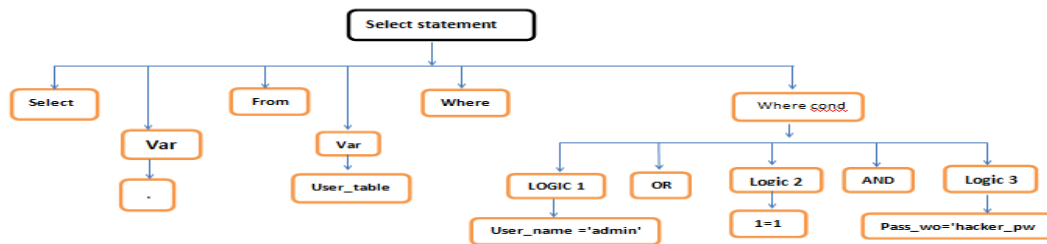


Figure.3. *SQL_Statement_Injection*

3.3 Related Work

The techniques related to SQL injection are classified and evaluated by [9]. In this section, we list the work related to ours and discuss their pros and cons. We can classify the related work to two main parts:

3.3.1 Detections SQL injection techniques

Vulnerability detection is an approach for detecting vulnerabilities in web applications, especially in the development and debugging phases. This approach is conducted either manually by developers or automatically with the use of vulnerability scanners. In the manual approach, an auditor manually reviews source code and/or attempts to execute real attacks to the web application. For discovering vulnerabilities, the auditor is required to be familiar with the software architecture and source code, and/or to be a computer security expert to attempt effective attacks tailored to his or her web application. A comprehensive audit requires a lot of time and its success depends entirely on the skill of the auditor. In addition, manual check is prone to mistakes and oversights. On the other hand, the vulnerability scanners automate the process of vulnerability detection without requiring the auditor to have detailed knowledge of the web applications including security details. The automated vulnerability scanners eliminate mistakes and oversights that is typically prone to be made by manual vulnerability detection.

From this reason, vulnerability scanners are widely used for detecting vulnerabilities in web applications.

The dynamic analysis scanners are based on penetration test, which evaluates the security of web applications by simulating an attack from a malicious user. The attack is typically generated by embedding an attack code into an innocent HTTP request. After sending the attack to the target web application, the vulnerability scanner captures the web application output to analyze the existence of vulnerabilities. Existing vulnerability scanners [10, 11, 12, 13, 14] employ dynamic analysis techniques for detecting vulnerabilities. AMNESIA combines static analysis and runtime monitoring [15]. In static phase, it builds models of the different types of queries which an application can legally generate at each point of access to the database. Machine Learning Approach Valeur [16] proposed the use of an intrusion detection system (IDS) based on a machine learning technique. IDS is trained using a set of typical application queries, builds models of the typical queries, and then monitors the application at runtime to identify the queries that do not match the model. The overall IDS quality depends on the quality of the training set; a poor training set would result in a large number of false positives and negatives. WAVES [17] is also based on a machine learning technique. WAVES is a web crawler that identifies vulnerable spots, and then builds attacks that target those spots based on a list of patterns and attack techniques. WAVES monitors the response from the application and uses a machine learning technique to improve the attack methodology. WAVES is better than traditional penetration testing, because it improves the attack methodology, but it cannot thoroughly check all the vulnerable spots like the traditional penetration testing. Instruction-Set Randomization SQLrand [18] provides a framework that allows developers to create SQL queries using randomized keywords instead of the normal SQL keywords. A proxy between the web application and the database intercepts SQL queries and de-randomizes the keywords. The SQL keywords injected by an attacker would not have been constructed by the randomized keywords, and thus the injected commands would result in a syntactically incorrect query. Since SQLrand uses a secret key to modify keywords, its security relies on attackers not being able to discover this key. SQLrand requires the application developer to rewrite code. SANIA [19] for detecting SQL injection vulnerabilities in web applications during the development and debugging phases. Sania intercepts the SQL queries between a web application and a database, and automatically generates elaborate attacks according to the syntax and semantics of the potentially vulnerable spots in the SQL queries. In addition, Sania compares the parse trees of the intended SQL query and those resulting after an attack to assess the safety of these spots. Paros is used for web application security assessment. Paros is written in Java, and people generally used this tool to evaluate the security of their web sites and the applications that they provide on web site. It is free of charge, and using Paros's you can exploit and modified all HTTP and HTTPS data among client and server along with form fields and cookies. In brief the functionality of scanner is as below. According to web site hierarchy server get scan, it checks for server miscount figuration. They add this feature because some URL paths can't be recognized and found by the crawler.

3.3.2 Preventing SQL injection techniques

Framework Support Recent frameworks for web applications provide a functionality that can be used to prevent SQL injections. For example, Struts[20] supports a validator. A validator verifies an input from the user conforms to the pre-defined format of each parameter. If a validator prohibits an input from including meta-characters, we can avoid SQL injections. Since a validator does not transform the dangerous characters to safe ones, we cannot prevent SQL injections if we want to include meta-characters in the input. Prepare Statement SQL provides the prepare statement, which separates the values in a query from the structure of SQL. The programmer defines a skeleton of an SQL query and then fills in the holes of the skeleton at runtime. The prepare statement makes it harder to inject SQL queries because the SQL structure cannot be changed. Hibernate [21] enforces us to use the prepare statement. To use the prepare statement, we must modify the web application entirely; all the legacy web applications must be re-written to reduce the possibility of SQL injections. Queries are intercepted before they are sent to the database and are checked against the statically built models, in dynamic phase. Queries that violate the model are prevented from accessing to the database. The primary limitation of this tool is that its success is dependent on the accuracy of its static analysis for building query models. In SQL Check [22] and SQL Guard [23] queries are checked at runtime based on a model which is expressed as a grammar that only accepts legal queries. SQL Guard examines the structure of the query before and after the addition of user-input based on the model. In SQL Check, the model is specified independently by the developer. Both approaches use a secret key to delimit user input during parsing by the runtime checker, so security of the approach is dependent on attackers not being able to discover the key. In two approaches developer should to modify code to use a special intermediate library or manually insert special markers into the code where user input is added to a dynamically generated query. CANDID [24] modifies web applications written in Java through a program transformation. This tool dynamically mines the programmer-intended query structure on any input and detects attacks by comparing it against the structure of the actual query issued. CANDID's natural and simple approach turns out to be very powerful for detection of SQL injection attacks.

We believe high precision can be achieved by discovering more vulnerabilities and by avoiding the issue of potentially useless attacks that can never be successful, with conducting fewer attacks. To achieve this, we focus on the technique of generating attacks that precisely exploit vulnerabilities. As a result, our approach generates attacks only necessary for identifying vulnerabilities. We propose SQLI-DP technique for detecting SQL injection vulnerabilities. In the output of the web application, which results in making fewer attacks, detecting more vulnerabilities, and making fewer false positives/negatives.

4. Problem Solution

4.1 Proposed technique for SQL Injection Detection and Prevention (SQLI-DP)

In this section, we present SQLI-DP technique, which tests for SQL injection vulnerabilities and prevent SQL injection by the parse of the queries. The general view of SQLI-DP is shown in figure 4 the discrete line if the software developer doesn't active SQLI-DP technique.

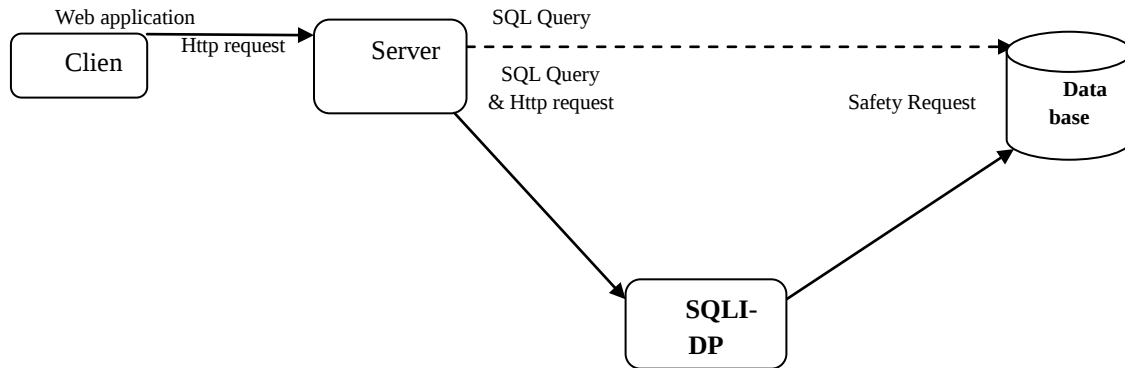


Figure4. General view of SQLI-DP technique

4.1.1 The main contribution of SQLI-DP technique:

1. Detect blind SQL injection using new function.
2. Using parse tree to detect SQL injection using Zql [25] with open source technique.
3. Detect and prevent the SQL injection using this leads to decrease execution time.

Our technique has two options (continue with safety, unsafe option) if we select the safety option activate SQLI-DP. If no (unsafe option) send request immediately to database and execute query. Illustrates the core work of SQLI-DP technique where the three points.

4.1.2 SQLI-DP Technique pseudo code

SQLI-DP technique (SQLIA DETETION & PREVENTION)

```

1. INPUT: SQL, Http request
2. OUTPUT: prevent attack & Final Report.
3. T ← Http request.
4. Q ← sql statement query.
5. BEGIN
6. IF (DETECTION BLIND SQLIA FOREACH T ) THEN
7.   {
8.     PREVENT T
9.     CREAT FINAL REPORT
10.    QUIT : SQLI-DP TECHNIQUE
11.  }
12. ELSE IF (DETECTION SQLIA FOREACH Q)
13.   {
14.     PREVENT T
15.     CREAT FINAL REPORT
16.     QUIT : SQLI-DP TECHNIQUE
17.   }
18.
  
```

```

19.  ELSE
20.  {
21.      SEND Q TO DATABASE
22.  }

```

4.1.3 The explanation steps of SQLI-DP technique is presents in details as follows

STEP 1:

Read the HTTP request from server and check if there is blind SQL injection or not as illustrated in [26] some example about blind SQL injection attack we use this method to detect blind SQLIA. There are several uses for the Blind SQL Injection:

- Testing the vulnerability.
- Finding the table name.
- Exporting a value.

There are more examples to traditional blind SQL injection and advanced blind SQL injection [26]. We noticed that from last example there are traditional kind of blind SQL injection and advanced kind with regular expression. To testing blind sql injection we scan the http request after capture if it contain the key words [select, top, version, user, order, substring, ascii, from, limit, having and etc.] or regular expression like [*,>,<,\$,%,-, '[a-z],[A-z],'" and etc] . The normal http request not contain like this. Then classify the http request to normal or abnormal request. If the system found Blind sql injection reject input go to step 3, if no continue to step 2.

STEP 2:

Perform SQL validation using validation parse tree. Compare the parse tree generated from an attack request with that generated from an innocent message to verify whether an attack was successful or not. If the parse generated from an innocent request, SQLI-DP determines the attack was successful. Suppose that a web application issues the SQL query "SELECT * FROM users WHERE name=""(vulnerable spot), or '1'='1".we use the technique idea that used in [27] when using a fresh key to user input. If the web application sanitizes the input properly, the parse tree will look like the one shown in figure 1. If not properly sanitized, the parse tree will look like the one shown in figure 3, where the structure of the parse tree is different from figure 1, if the two tree are different detect there is injection. Then classify the query as normal or abnormal and account the False positive (A false positive occurs when the test returns a positive result, but there is actually no fault). Then Prevent abnormal user queries and send normal user queries immediately to database and execute query and send the data to server.

STEP 3:

SQLI-DP generates the report that contains the SQL query with the user input data, if it found blind sql injection and execution time by millisecond. The SQL query benefit in detection to know the kind of SQL injection attack.

The advantages of SQLI-DP are illustrated as follows:

- 1- It is efficient, because it only adds about 1second overhead to database query costs.
- 2- In addition, it is easily adopted by software developer.
- 3- It is suitable for legacy system because it is a technique implemented on server side. It doesn't need to rewrite old web application because protection from Injection is on server side where database resides.

4.2 Experimental Result

The experiments methodology is done by designing MSQL database and a patche server web application in our platform. The SQLI-DP technique is implemented by Java language. We used Zql to implement an SQL parser. The SQLI-DP technique consists of an SQL proxy, and a core component of our technique. The SQL proxy captures the SQL queries. The core component of SQLI-DP performs the tasks described in the previous section. We apply SQLI-DPs in windows7, 32-bit operating system and core (TM) i5 CPU. Some snapshot screens from SQLI-DP work are shown when using different options of unsafe and safe. If we using unsafe system with SqlIA we noticed that the attack success and the server return all data records in database table as shown in figure 5. SQLI-DP (unsafe option) response all recodes because there are (Tautology injection found). If we using Safety system option SQLI-DP prevents SQLIA in first test if it found blind injection attack it generates report as shown in figure 6. If blind injection attack not found the SQLI-DP tests other types of sql injection attacks. Figures 7 and 8 are show the reports generated after hacking in different types of attack.

This report is benefit in:

- 1- To know the type of SQL injection attack.
- 2- Sure that the system prevents this attack to accessed database.
- 3- To estimate the execution time of detection and prevention.

As shown in figure 7 this report appears when the SQLI-DP found **Piggy-Backed Queries** SQLIAs when try to access the web application database.

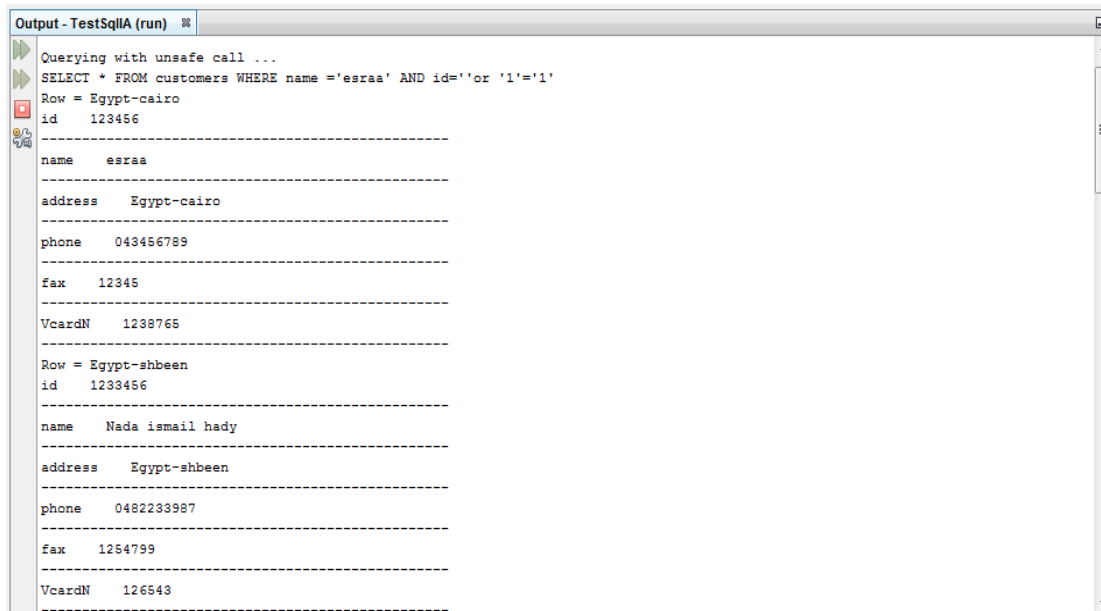


Figure 5. Using unsafe option



Figure 6. SQLI-DP safety option with Blind sql injection

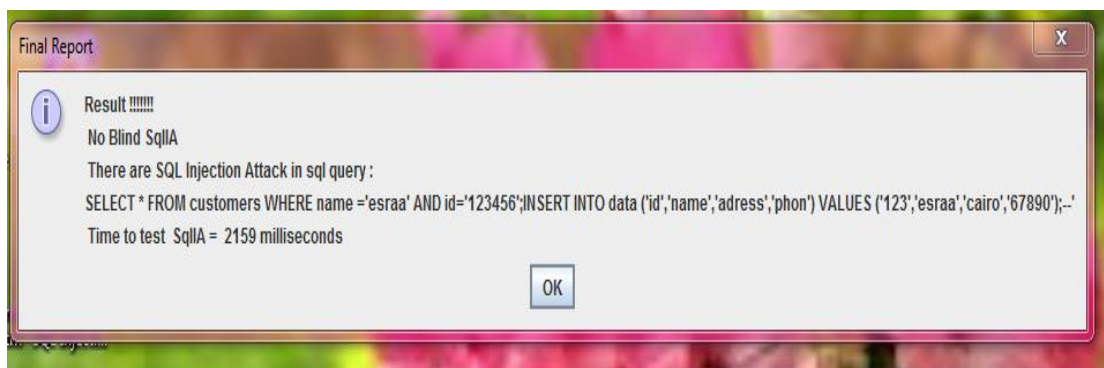


Figure 7. SQLI-DP with safety option with Piggy-Backed Queries SqlIAs

As shown in figure 8 this report appears when the SQLI-DP find Tautology SqlIAs when try to access the web application database.

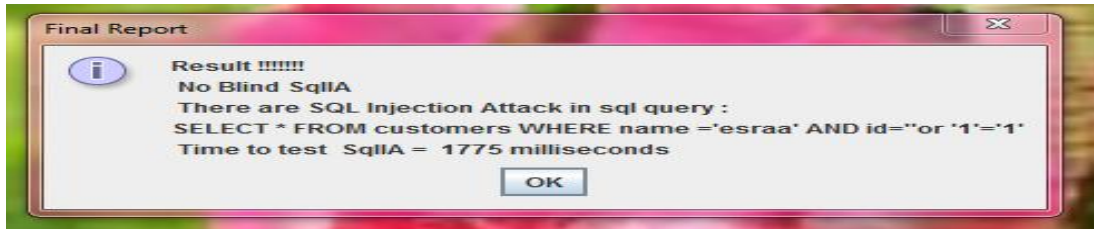


Figure8. SQLI-DP with safety option with Tautology SqlIAs.

Table1 shows the SQLI-DP techniques with respect to the number of attacks trails we using manual hacking on the database and the execution time in millisecond.

Table 2. SQLI-DP trails and execution time

Types of SQLIA	Num.of Trials of attacks	Detection & prevention	Execution time(ms)
Tautology	50	(50) 100%	1775
Logically Incorrect Queries	50	(50) 100%	1670
Union Query	50	(50) 100%	1819
<i>Piggy-Backed Queries</i>	50	(50) 100%	1670
<i>Blind injection</i>	40	(40) 100%	102
<i>Alternate Encoding</i>	10	(10) 100%	1560
<i>Stored Procedure</i>	10	(10) 100%	1800

We compare our SQLI-DP technique with Paros (detection technique).

- During our work we observed The SQLI-DP technique detect and prevent the Blind SQLIA as the first check.
- The SQLI-DP technique give the user request query but the Paros technique no.
- The SQLI-DP technique add about 1s addition to the system but Paros technique add 5 second.
- The SQLI-DP technique is Server Side protection technique but Paros technique host side protection.
- As shown in figure 9 the execution time to detect attack in SQLI-DP technique less than Paros technique.

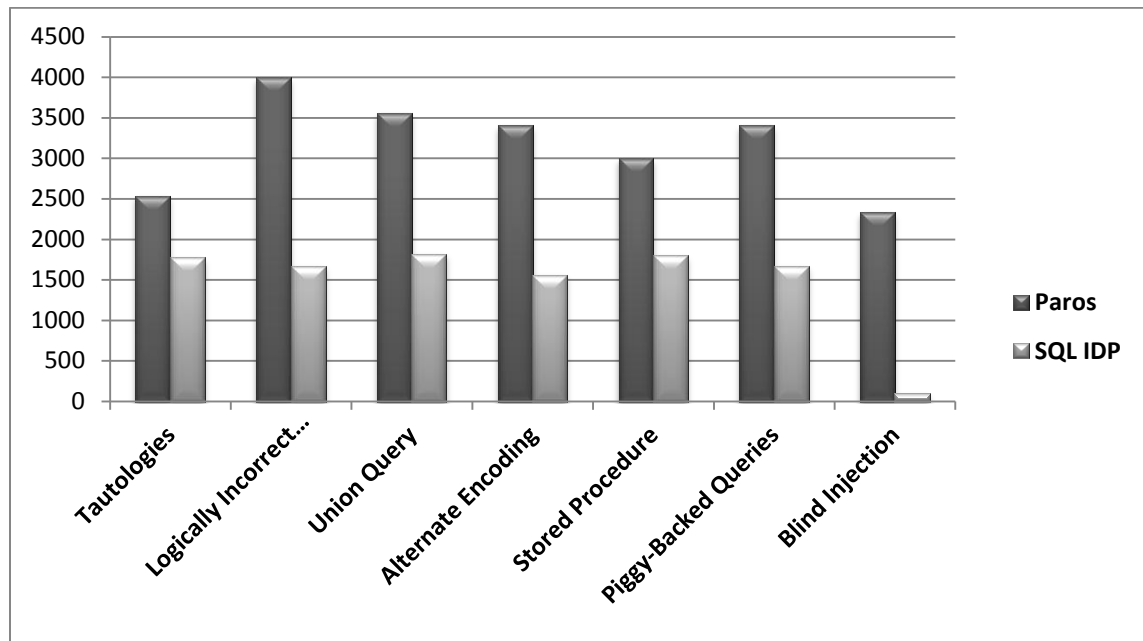


Figure9. Execution time of SQLI-DP and Paros techniques

Through the study of previous observations and discuss the reasons, we find that SQLI-DP by several features not found in other techniques. First, SQLI-DP detects and prevents using less time than Paros for the following reasons:

- 1- SQLI-DP is server side technique work as a proxy between the server and database but Paros work in host side is a proxy between client and server so SQLI-DP decreases the communication time.
- 2- SQLI-DP detects and prevents the blind SQL injection attacks but Paros technique detect only.
- 3- From the observation the SQLI-DP technique add only one second addition to the system but Paros technique add approximately 4 second to the system.
- 4- The SQLI-DP technique gives the SQL query statement in the final report it's very important to the developer:
 - To know kind of attack.
 - To detect the kind of attacks and use for protect this type of sits from this kind of attacks.
 - To know the advanced schema of SQL injection attacks.
- 5- SQLI-DP give the final report after prevent the attacks to sure the attacks are prevented.
- 6- In addition, the SQLI-DP it is easily adopted by software developer because it written by Java language that suitable for any platform .
- 7- It is suitable for legacy system because it is a technique that implemented on server side. It doesn't need to rewrite old web application because protection from Injection is on server side where database resides.

5. Conclusion and Future Work

We presented our new server side protection technique against SQL Injection, which is designed to check for SQL injection vulnerabilities in the server side. SQLI-DP intercepts SQL queries and analyze it based on the syntax of the SQL queries using parse tree. It rejects injectable queries from beginning and executes other queries to detect SQL injection. It is suitable for legacy system because it is a technique that is implemented on server side. It doesn't need to rewrite old web application because protection from injection is on server side where database resides. SQLI-DP has two advantages comparing with other scanner. It's efficient, adding about 1second overhead to database query costs. In addition, it is easily adopted by software developer, having the same syntactic structure as current popular record set retrieval method and we detect and prevent blind SQL injection.

In future work, we intend to evaluate SQLI-DPs using different web based applications with real public domain to achieve great accuracy in SQL injection detection and prevention.

6. References

- [1] OWASP (2012, September 6). Testing for SQL Injection[Online]. Available: [http://www.owasp.org/index.php/Testing for SQL Injection](http://www.owasp.org/index.php/Testing_for_SQL_Injection).
- [2] Nithya. A Survey on SQL Injection attacks, their Detection and Prevention Techniques. International Journal Of Engineering And Computer Science Volume 2 Issue 4 April,2013 Page No.886-905.
- [3] Sayyed Mohammad. Study of SQL Injection Attacks and Countermeasures. International Journal of Computer and Communication Engineering, Vol. 2, No. 5, September 2013 Page No.539-514.
- [4] Gopi Krishnan. Preventing Injection Attack by Whitelisting Inputs . Lecture Notes on Information Theory Vol. 1, No. 3, September 2013 Page No.132-134.
- [5] W. Halfond and A. Orso, "Combining Static Analysis and Runtime Monitoring to Counter SQL- Injection Attacks," Proceeding of the Third International ICSE Workshop on Dynamic Analysis (WODA 2005), 2005.
- [6] Chunhui Song. Song, "SQL Injection Attacks and Countermeasures", California State University, Sacramento, (Spring 2010).
- [7] Z. Lashkaripour .A Simple and Fast Technique for Detection and Prevention of SQL Injection Attacks. International Journal of Security and Its Applications Vol.7, No.5 (2013), Page 53-66.
- [8] K.R Venugopal, L.M Patnaik , Patnaik,"Detection and prevention of SQL injection attacks", Computer Networks and Intelligent Computing. 5th International Conference on Information Processing, ICIP 2011, Bangalore, India, August 5-7, 2011, pages 104-107.

- [9] Atefeh Tajpour, Maslin Masrom, Mohammad Zaman Heydari, Suhaimi Ibrahim,"Evaluation of SQL Injection Detection and Prevention Techniques," 2nd International Conference on Computational Intelligence, Communication Systems and Networks, Liverpool, United Kingdom pages 216-221.
- [10] Jason Sabin & Dan Timpson (2006, August). Paros (version 3.2.13) [Online]. Available: http://www.testingsecurity.com/paros_proxy.
- [11] Acunetix. Acunetix Web Security Scanner. <http://www.acunetix.com/>.
- [12] IBM. Rational AppScan. <http://www.ibm.com/software/awdtools/appscan/>.
- [13] Yao-Wen Huang, Shih-Kun Huang, Tsung-Po Lin, and Chung-Hung Tsai. Web Application Security Assessment by Fault Injection and Behavior Monitoring. In Proceedings of the 12th International Conference on World Wide Web, May 2003, pages 148–159.
- [14] Stefan Kals, Engin Kirda, Christopher Kruegel, and Nenad Jovanovic. SecuBat: A Web Vulnerability Scanner. In Proceedings of the 15th International Conference on World Wide Web, May 2006, pages. 247–256.
- [15] W. Halfond and A. Orso. AMNESIA: Analysis and Monitoring for NEutralizing SQL-Injection Attacks. In Proceedings of the 20th IEEE/ACM International Conference on Automated Software Engineering (ASE), pages 174–183, 2005.
- [16] F. Valeur, D. Mutz, and G. Vigna. A Learning-Based Approach to the Detection of SQL Attacks. In Proceedings of the Conference on Detection of Intrusions and Malware and Vulnerability Assessment (DIMVA), pages 123–140, 2005.
- [17] Y. Huang, S. Huang, T. Lin, and C. Tsai. Web Application Security Assessment by Fault Injection and Behavior Monitoring . In Proceedings of the 12th International World Wide Web Conference (WWW03), pages 148–159, 2003.
- [18] S. Boyd and A. Keromytis. SQLrand: Preventing SQL injection attacks. In Proceedings of the Applied Cryptography and Network Security (ACNS), pages 292–304, 2004.
- [19] Yuji Kosuga, Kenji Kono, Miyuki Hanaoka, (August 2011). Sania: Syntactic and Semantic Analysis for Automated Testing against SQL Injection, Doctor of Philosophy. Keio University.
- [20] Struts. Apache Struts project. <http://struts.apache.org/>.
- [21] Hibernate. hibernate.org. <http://www.hibernate.org/>.
- [22] Z. Su and G. Wassermann. The Essence of Command Injection Attacks in Web Applications. ACM SIGPLAN Notices. Volume: 41, pp: 372-382, 2006.
- [23] Gregory T. Buehrer, Bruce W. Weide, and Paolo A. G. Sivilotti Using : Parse Tree Validation to Prevent SQL Injection Attacks, Computer Science and Engineering The Ohio State University, pages 4-5, 2009.
- [24] Sruthi Bandhakavi, Prithvi Bisht, P. Madhusudan, CANDID: Preventing SQL Injection Attacks using Dynamic Candidate Evaluation. Proceedings of the 14th ACM conference on Computer and communications security. ACM, Alexandria, Virginia, USA. page: 12-24.
- [25] Gibello. Zql: A java sql parser, 2002. In <http://www.experlog.com/gibello/zql/>.

- [26] Simone 'R00T_ATI' Quatrini Marco 'white_sheep' Rondini. Blind Sql Injection with Regular Expressions Attack.
- [27] Jagdish Halde. (2008),SQL Injection analysis, Detection and Prevention. Master's Theses and Graduate ResearchSan Jose State University.

3D Object Categorization Using Spin-Images with MPI Parallel Implementation

A. Eleliemy
ahmed_hamdy@
cis.asu.edu.eg

D. Hegazy
Computer Systems Department
Ain Shams University
Doaa.hegazy@cis.asu.edu.eg

W. Elkilani
Wail.elkilani@cis.asu.edu.eg

Abstract: Object recognition and categorization are two important key features of computer vision. Accuracy aspects represent research challenge for both object recognition and categorization techniques. High performance computing (HPC) technologies usually manage the increasing time and complexity of computations. In this paper, an approach utilizing 3D spin-images for 3D object categorization is proposed. The main contribution of this approach is that it employs the MPI techniques in a unique way to extract spin-images. The technique proposed utilizes the independence between spin-images generated at each point. Time estimation of this technique has shown dramatic decrease of the categorization time proportion to the number of workers used.

Keywords: Object categorization, spin-image, support vector machine, MPI.

17. Introduction

3D object categorization is one of the important fields in computer vision. Its goal is making the computer (e.g. robot) able to recognize general classes of object categories, which is more human-like ability. 3D object categorization has many different applications including robot navigation and surveillance, security systems and manufacturing quality control. Such variety of applications lead to a large different sets of constraints which give each problem its flavors. In recent few years, the growing demand for on-line categorization systems has motivated a lot of researchers on this scope. So the majority of this work will target the on-line categorization systems using HPC techniques. The model proposed uses spin-image features [1] for recognizing and categorizing 3D objects. Spin images are considered an important 3D feature, used by different approaches for 3D object recognition task [1-5]. However, their high computational performance is its main disadvantage.

The proposed HPC model for 3D object categorization focuses on using MPI technique to accelerate the spin-images generation process. Moreover spin-images generated can be used for feature training of Support Vector Machine (SVM) classifier instead of recognition by surface matching. This paper assumes that the optimum system is subject to the following three constraints:

- *3D object nature constraint:* This constraint is inherited from computer vision goal where the goal is to make computers see like humans. Humans can realize the third dimension but digital images didn't have such information except in the last few years when range cameras and depth sensors have great development. While techniques needed to infer properties of the 3D world from digital images, these devices give the depth information for the scene. Although the third dimension could be obtained by stereo images, yet development in range and depth sensor technology allows easier use

of such information [6]. The third dimension information adds more reality for the problem domain. So the system must deal with 3D world information.

- *Real time Constraint:* What is meant by real time constraint is the system ability to process at least from 10 frame per second (fps) up to 33 (fps) [7]. Such variety comes from the system domain. In many techniques the timing constraint is tied to the categorization success rate. To increase success categorization, it may increase the technique complexity which means more time. So categorization success rate will be treated as a separate constraint.
- *Success categorization rate:* This constraint is to show how system could be tolerant. Depending on application some categorization techniques could allow false categorization rate more than others.

In this paper, we present an approach for 3D object categorization using spin-image features and Support Vector Machines (SVM) classifier with a MPI parallel implementation. The proposed model tries to meet the mentioned constraints to achieve an optimum categorization performance.

The paper is organized as follows: Section II presents the related work. The proposed 3D categorization model and its parallel version are discussed in section III. Section IV presents the experimental evaluation and the conclusions are given in section V.

18. Related Work

Due to the importance of 3D object categorization, many researchers have targeted the problem which resulted in different categorization approaches such as (8). However, only few approaches addressed the on-line categorization using HPC.

As we mentioned before, spin-images are used by different categorization approaches. One of these approaches is presented in [1], where spin-images are used to recognize multiple

3D objects in a scene for 3D object retrieval. This approach is based on surface matching, by matching surfaces of already stored objects with the new two incoming captured objects. Spin-images for each known model are generated and stored. In the recognition step, spin-images are generated for each point of the incoming captured scene. Finally the similarity between the calculated new points and the spin-images from all stored models is measured. Best matching means both of recognition and localization for the object are done. Although such work doesn't discuss any timing measures, but it is clear that such technique consists of two heavy computational tasks: Spin-images Generation process and Surface Matching. Such research [1] has been extended and improved in [9]. Instead of using the original spin-image algorithm, they introduced an enhanced version called spin-image signatures. Spin image signatures used for content-based retrieval of 3D Objects and shows some important results in terms of sensitivity to model deformations. Such approach includes three steps: 1-original spin-images generation, 2-spin image signatures, 3- clustering of spin-image signatures. Total complexity is in order of $O(n^4)$, where n represents the vertices' of 3D object. Such approach violates the time constraint. In [10], the photometric information is considered. New version of the original spin-image algorithm, called textured spin-images, is introduced. Textured spin-images are used in

3D registration. Such algorithm enjoys remarkable properties, since it can give rigid motion estimates, more robust, more precise, and more resilient to noise than standard spin-images. But this approach involves the same steps to obtain spin-images of 3D object. Although it reduces complexity, it still involves intensive computation for spin-image generation and calculation.

Also spin-image algorithm is a base for some research at the area of face recognition. In [4], a feature based method for face recognition is introduced. Surface shape information is inferred from image brightness using a Lambertian shape- from-shading scheme. Spin-image was used to perform surface representation from the surface normal. To reduce the computation of spin-images, local spin-images are computed on the image patches using surface normal information. Although such approach could avoid some of the spin-image algorithm complexity, it calculates spin-images at some points not at each vertex which means the spin-image generation process is still as it is.

In [11], has introduced a recognition system based on a hardware accelerated implementation using Field Programmable Gate Array (FPGA). The proposed hardware is used in spin-images generation process and surface matching have. Such work enhances the overall system run time from several minutes in [1] to 4.5 seconds. However, FPGA is not the only possibility for performance enhancement different possibilities to give better performance, such as pipelining, lookup-tables for square roots and input-output stage.

Another approach is presented in [12], where the heavy task of spin-image surface matching is passed. This work introduces 3D object categorization based on spin-images and bag of features classification instead of surface matching. Such approach gives significant results in terms of recognition rate 62.5% at 100 spin-images per object. But the problem of heavy spin-images generation calculation is not avoided. There is no concern for the overall system runtime. Therefore the approach presented in [12] violates real-time constrains.

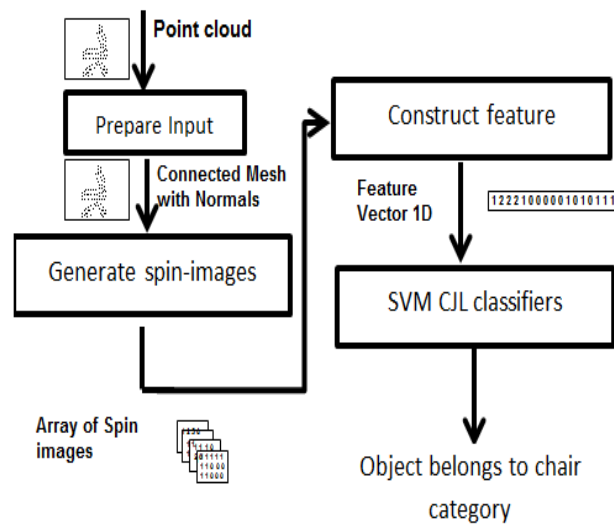


Fig. 1. Sequential 3D object categorization system

It can be noted from the previous discussion that the main drawback of the spin-image algorithms proposed is the heavy computation involved. Most authors have concentrated on elevating the recognition rate without considering system run time. Moreover, as far we have seen that neither of authors have parallelized spin-image generation algorithm using MPI. In the proposed approach, we have utilized the MPI technology to reduce dramatically the overall system runtime with maintaining the recognition rate. In fact we have enhanced slightly the recognition rate.

19. Proposed Categorization Model

A. Sequential 3D object categorization model

The proposed 3D object categorization model consists of four main steps, as shown in Figure 1. The first is the preparation step. Normally, 3D objects are obtained from different depth-image sensors in the form of point cloud. Such input needs to be converted into the form of connected mesh. Connected meshes could be obtained from point cloud by many triangulation techniques such surface reconstruction from unorganized points in [13], 3D mesh reconstruction from point cloud using elementary Vector and geometry analysis in [14] or reconstruction from point clouds in 3D in [15]. However the data gathered is synthetic 3D objects not a sensory input images as in mentioned in section IV. This means that the data is already in the form of triangulated mesh. Therefore, in the preparation step in the built model the normal vector for each point in the mesh surface is calculated. In the second step of the model, spin-images for each object are generated. Then, the generated spin-images are converted to feature vector. So, each object is represented in the model by its corresponding feature vector. Finally, the calculated feature vectors are learned and classified using SVM classifier. Now, we will give a brief description of the spin-image generation as it is an important step in the proposed model.

Calculating Spin-images: Spin-image is a data level shape descriptor used for recognition of multiple objects. Spin- images are introduced in [1] in manner that recognition is done based on matching surfaces by matching points using the spin-image representation. Spin-image is calculated for each oriented point in the 3D mesh. Figure 5 shows the generation process of spin-images at point P. According to equations to1and 2, two values (α , β) are calculated with each oriented point like Point X.

$$\alpha = \sqrt{||x - p||^2 - (n \cdot (x - p))^2} \quad (1)$$

$$\beta = n \cdot (x - p) \quad (2)$$

In [1] the analogy to describe the image generation process, a white paper which is rotated around the normal at the point, when they touch the paper, they stick to it and form a 2D cumulative image. The generation of spin-images is affected by four important parameters: bin size, support angle, image width and support distance. Values of these parameters depend on the objects data set. Bin size determines the storage size and the descriptiveness of the spin images. The best value for bin- size is a multiple of mesh surface resolution (1). Image width

means number of rows or columns as square spin-image is used. Support angle is the maximum angle between normal of the oriented point and the normal of points that will participate in generated spin-image. Support distance equals image-width multiplied by bin-size. Estimating the best values of each parameter is done with experimental evolutions explained in section IV.

After the spin-images generation done. The resulted 2D spin-images for each 3D object are converted to be the feature vector of these object, by converting the 2D spin-image into 1D vector where all rows in 2D spin-image became at one row. The feature vectors are learned and classified using support vector machines (SVM). SVM classifiers are famous and widely used classifiers that showed robust performance in many classification tasks [16]. Due to space limitation, the reader can refer to [8] for more details. However, we provide here a short explanation of how new images are classified in the proposed approach. After training the SVM classifiers using the spin-image features, a list of SVM binary classifiers with length L is generated, where $L = nC_2$ represents the number of categories (six categories are mentioned in section IV). This is called the classifier: Chain Judges List (CJL). Figure (2) shows how judges in the (CJL) classifier work. Assume we have 4 categories W, X, Y and Z, so judges list consist of 6 judges $J_{WX}, J_{WY}, J_{WZ}, J_{XY}, J_{XZ}$ and J_{YZ} . J_{XY} means that this SVM binary classifier is trained with the spin-image features of the images of two categories X and Y. Training each classifier in the six classifiers is performed in independent manner, each represents an object category. After training, a new previously unseen image is presented to the system for testing and the spin-image feature is calculated. The test feature vector is then presented to each binary classifier. All classifiers vote for the input image and the classifier (category) with maximum votes will win.

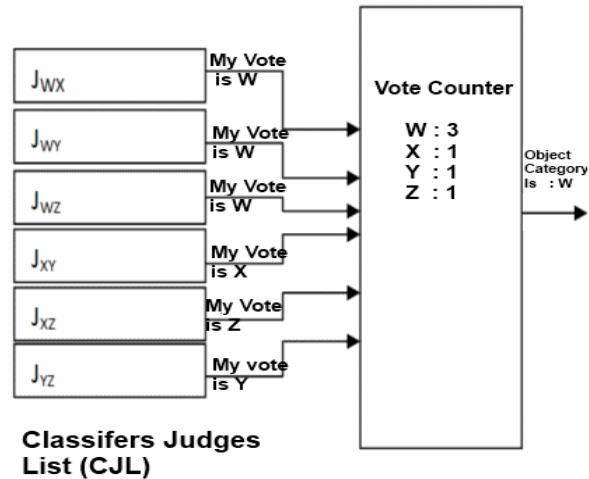


Fig. 2 Classifier judges list for different 4 categories W, X, Y and Z

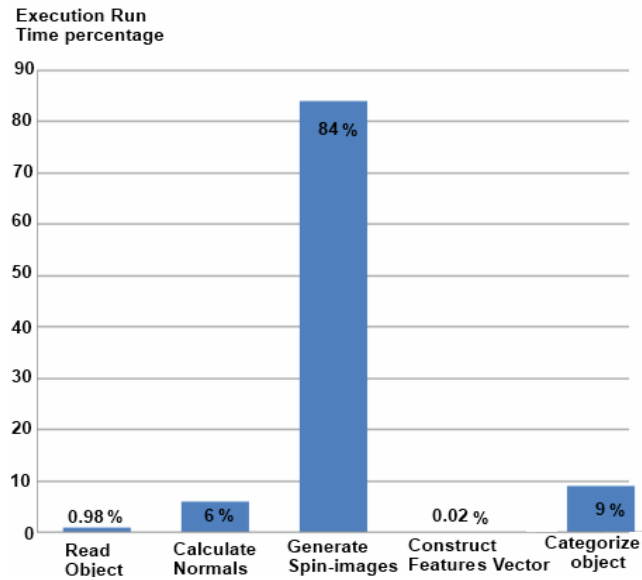


Fig. 3 System runtime analysis

B. 3Dparallelcategorizationmodel

In the sequential model the spin-image generation process consumes the majority of system time (almost 84% of system time). Figure 3 records how the system time is consumed by different system stages. According to equation 1, the task dependency between generations of spin-images is minimal. The idea is to convert the sequential spin-images generation process to parallel form. Figure 4 shows the parallel categorization model. Based on message passing interface (MPI), a group of workers cooperate to perform the 3D object categorization process.

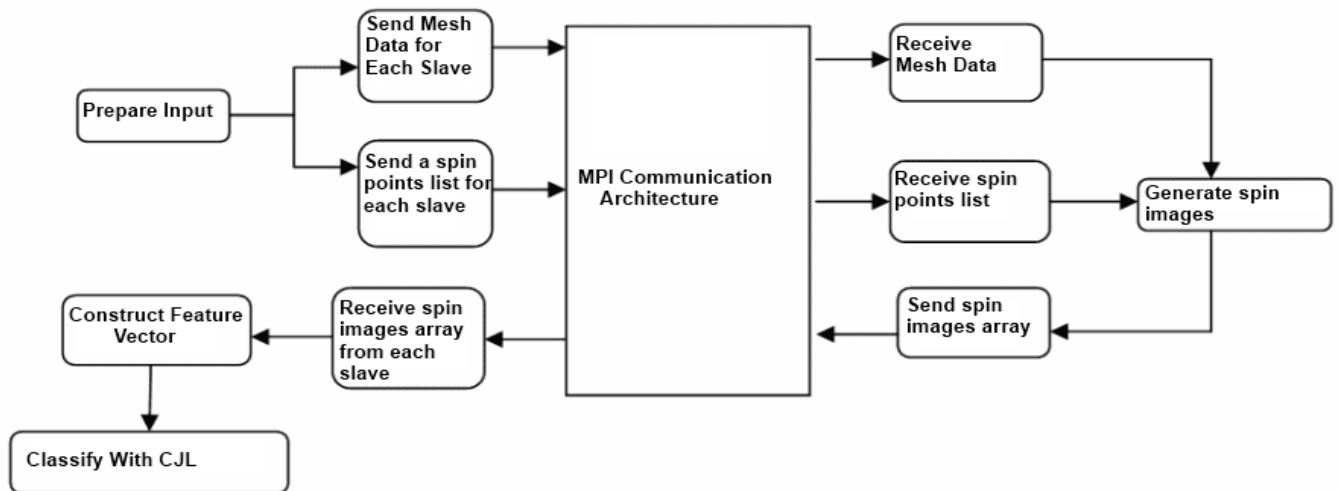


Fig. 4 Parallel 3D object categorization

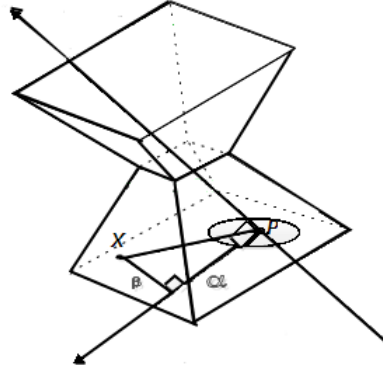


Fig. 5 The generation process of spin-image at Point P with normal n for 3D object

Based on Master/slaves model, one of these workers is the master who performs the preparation for the input 3D information. The master prepares the input to be in form of connect faces (Mesh) like in figure 5, also calculates the normal at each point. The master then sends both of prepared mesh data and list of spin-points to each slave. Each slave receives the mesh data and the list of spin-points and then calculates its spin-images. Slaves send the spin-images to the master worker which converts spin-images into feature vector (observation) to be classified according to the trained judges list as in the proposed sequential model. Category of maximum judges' votes is the decision.

20. Experimental Evaluation

A.Data

The Princeton Shape Benchmark (PSB) is used in all the experimental evaluation of the proposed model. PSB is one of the public available databases that contains a variety of synthetic3D models.

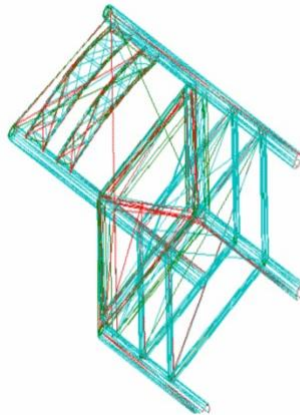


Fig. 6 Sample of Mesh object constructed at master worker

TABLE. I TABLE OF PSB FURNITURE CLASSES

Class	Number of objects
Table	69
Chair	29
Couch	15
Cabinet	9
shelf	26
Bed	8

Therefore the generated images don't suffer from noise, like in sensory captured data. Each class in PSB contains only one object [17]. The furniture objects are selected to be the focus of the proposed categorization model. Table I shows how many samples of each object are available. Also Figure 6 shows rendering of some PSB 3D Furniture objects. As mentioned

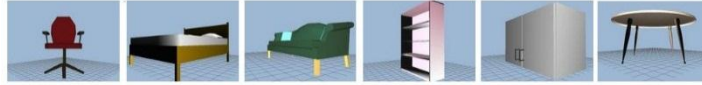


Fig. 7 Some of PSB furniture objects

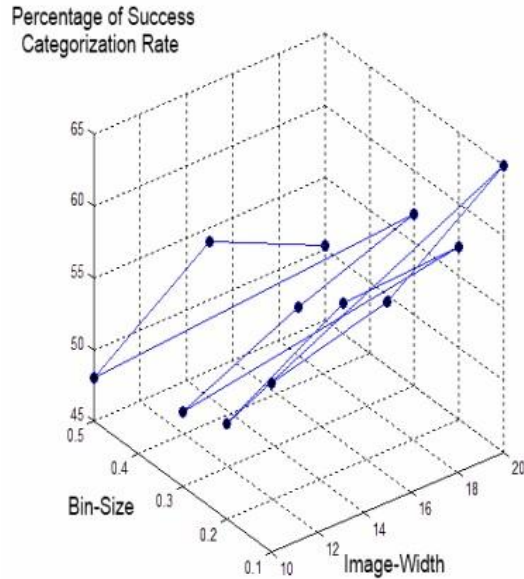


Fig. 8 Effect of bin-size parameter, support-angle and image-width on the mean of categorization rate

before, spin image generation process is done at only the oriented points. Objects in PSB are stored in a format which contains points and faces of object made by these points. We divide the gathered data into two parts: one for training and the other one is for testing. The ratio between the two parts is 80% and 20% respectively.

B.Results

As mentioned before, the quality of generated spin-images are related to the generation parameters, also they have direct effect on the categorization process. Several experiments have been already performed to select be the best parameters. Figure 8 shows the increase of bin-size for decrease of the success categorization rate. In the other side, the increase of the image width causes the increase of the success categorization rate. The best success categorization rate has been recorded at the intersection point where image-width is 20 pixels and bin-size equals 0.1.

In the proposed model, we didn't look up for minimum spin points to get highest categorization rate. Instead of that we tried different number of random spin point to identify the best categorization rate. As shown in Figure 9, the best success categorization rate has been recorded at 100 random spin points.

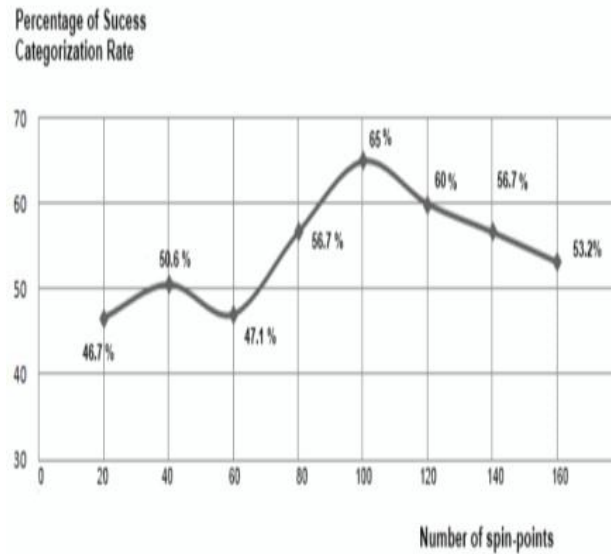


Fig. 9 Mean of success categorization rate at different number spin-point

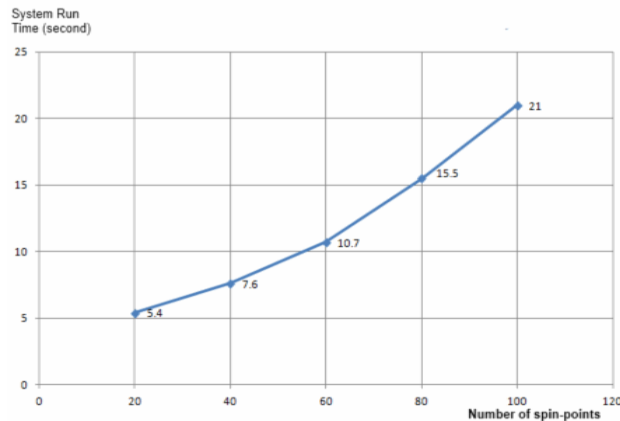


Fig. 10 Relation between system overall runtime and number of spin-points

Figure 10 shows the relation between system execution run time and number of spin points. While the number of taken spin points increases, the total amount of system execution time increases. As in figure 9, 100 spin point have best categorization rate which means 21 seconds as a total recognition time.

The confusion matrix in figure 11 shows how failure and success categorization is distributed among different system categories. For example the first row, 66% of cabinet object are correctly categorized to Cabinet, while 33% are incorrectly categorized to chair category.

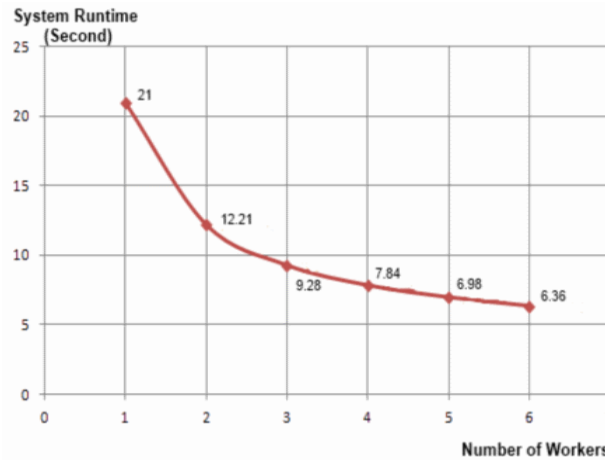
Ground Truth	Predicted					
	Cabinet	Bed	Couch	Chair	Shelf	Table
Cabine	0.66	0	0	0.33	0	0
Bed	.0667	.7333	.0667	.0667	.0667	0
Couch	0	0	1	0	0	0
Chair	0	0	0	1	0	0
Shelf	0	.3333	0	.3333	.3333	0
Table	0	.1667	.3333	.1667	.1667	.1667

Fig. 11 Confusion matrix of proposed model shows the mean of success categorization rate equal to 65%

TABLE 2. TABLE OF SEQUENTIAL AND PARALLEL PARTS OF THE PROPOSED APPROACH

Number of workers	Time of Sequential Part	Time of parallel part
1	3.36	17.64 s
2	3.36	8.85 s
3	3.36	5.92 s
4	3.36	4.48 s
5	3.36	3.62 s
6	3.36	3 s

Figure 12 shows the results of the parallelized system. Increasing the number of the workers in our approach will decreased the total amount of categorization time. The proposed parallel approach has two parts: Sequential part and parallel part. Table II shows how sequential part is independent on number of workers while the parallel part is directly dependent on the number of workers. For six workers, we have reach a recognition time of 6.36 seconds. As listed previously that [11] used parallelized FPGA approach has reached 4.5 seconds. The suggested MPI approach is preferable for the well-known merits of software parallelized approaches. Obviously, that increasing number of workers will minimize categorization time runtime until the number of workers equates the number of spin points. Driving empirical formula for their relationship will be interesting future research work.

**Fig. 12 Relation between number of workers and overall system run time**

21. Conclusion and Future Work

Although using spin-images with SVM classifiers in 3D object categorization includes a heavy processing task, we used a distributed parallel model based on MPI to speedup the spin-image generation process which consumes the majority of the model processing power. As we can drive the maximum speed up obtained by our parallel proposed model. We recommend

such approach especially for objects with high 3D point density. Also we can parallelize the generation of single spin-image at each worker. We are planning to switch the approach proposed to use hybrid model of MPI and Open MP. We intend to use the kinect camera as sensory input and depth- image sensory database like (18).

22. References

- [1] A. E. Johnson and M. Hebert, "Using spin images for efficient object recognition in cluttered 3d scenes," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, pp. 433–449, 1999.
- [2] L. Wang, J. Sun, and Q. Wang, "Modified spin- image target recognition algorithm applied to laser radar range imagery," *Conference on Lasers and Electro- Optics/Pacific Rim 2009*, Optical Society of America, 2009.
- [3] T. Tamaki, S. Tanigawa, Y. Ueno, B. Raychev, and K. Kaneda, "Scale matching of 3d point clouds by finding key scales with spin images," *2010 20th International Conference on Pattern Recognition (ICPR)*, pp. 3480–3483, 2010.
- [4] M. Bock, G. Cortelazzo, C. Ferrari, and C. Guerra, "Identifying similar surface patches on proteins using a spin-image surface representation," in *Combinatorial Pattern Matching*, ser. Lecture Notes in Computer Science, A. Apostolico, M. Crochemore, and K. Park, Eds. Springer Berlin Heidelberg, vol. 3537, pp. 417–428, 2005.
- [5] J. Assfalg, A. Del Bimbo, and P. Pala, "Spin images for retrieval of 3d objects by local and global similarity," *Proceedings of the 17th International Conference on Pattern Recognition (ICPR 2004)*, vol. 3, , pp. 906–909, 2004.
- [6] B. Langmann, K. Hartmann, and O. Löffel, "Depth camera technology comparison and performance evaluation," *ICPRAM (2)*, P. L. Carmona, J. S. Sánchez, and A. L. N. Fred, Eds, SciTePress, p. 438–444, 2012.
- [7] J. Oh, G. Kim, I. Hong, J. Park, S. Lee, J.-Y. Kim, J.-H. Woo, and H.-J. Yoo, "Low-power, real-time object-recognition processors for mobile vision systems," *Micro IEEE*, vol. 32, no. 6, pp. 38–50, 2012.
- [8] C. Cortes and V. Vapnik, "Support-vector networks", *Machine Learning*, springer, pp. 273–297, 1995.
- [9] K. Veropoulos, G. Bebis, and M. Webster, "Investigating the impact of face categorization on recognition performance," *Advances in Visual Computing, LNCS*, Vol. 3804, pp 207–218, 2005
- [10] N. Brusco, M. Andreetto, A. Giorgi, and G. Cortelazzo, "3d registration by textured spin-images," *Fifth International Conference on 3-D Digital Imaging and Modeling (3DIM 2005)*, pp. 262–269, 2005
- [11] J. Nylander and H. Wallin, "Improving the performance of spin-image algorithm hardware acceleration in vhdl," Master's thesis, Malardalen university, Vasteras department of computer science and electronics, 2008.
- [12] M. Baum, D. Agnes, and D. Sven, "Using spin images for 3d object classification," Master's thesis, Applied Informatics Group, Bielefeld, 2011.

- [13] H. Hoppe, T. DeRose, T. Duchamp, J. McDonald, and W. Stuetzle, "Surface reconstruction from unorganized points," SIGGRAPH Comput. Graph. vol. 26, no. 2, pp. 71–78, Jul. 1992
- [14] N. Wongwaen, S. Tiendee, and C. Sinthanayothin, "Method of 3d mesh reconstruction from point cloud using elementary vector and geometry analysis," 8th International Conference on Information Science and Digital Content Technology (ICIDT), vol. 1, pp. 156–159, 2012.
- [15] L. Alboul and G. Chliveros, "A system for reconstruction from point clouds in 3d: Simplification and mesh representation," 2010 11th International Conference on Control Automation Robotics Vision (ICARCV), pp. 2301–2306, 2010.
- [16] H. Song, K. K. Choi, I. Lee, L. Zhao, and D. Lamb, "Adaptive virtual support vector machine for reliability analysis of high-dimensional problems." Structural and Multidisciplinary Optimization, vol. 47, pp. 479–491, 2013.
- [17] P. Shilane, P. Min, M. Kazhdan, and T. Funkhouser, "The Princeton shape benchmark," Internationalin Shape Modeling, Jun. 2004.
- [18] K. Lai, L. Bo, X. Ren, and D. Fox, "A large-scale hierarchical multiview rgb-d object dataset," 2011 IEEE International Conference on Robotics and Automation (ICRA), pp. 1817–1824, 2011.