



AFRL-RH-WP-TR-2011-0019

**AN INTELLIGENT DECISION SUPPORT SYSTEM FOR
WORKFORCE FORECAST**

**Hung-da Wan
Fengshan F. Chen
Glenn W. Kuriger**

**The University of Texas at San Antonio
One UTSA Circle
San Antonio TX 78249**

**JANUARY, 2011
Final Report**

Distribution A: Approved for public release; distribution is unlimited.

**AIR FORCE RESEARCH LABORATORY
711TH HUMAN PERFORMANCE WING,
HUMAN EFFECTIVENESS DIRECTORATE,
WRIGHT-PATTERSON AIR FORCE BASE, OH 45433
AIR FORCE MATERIEL COMMAND
UNITED STATES AIR FORCE**

NOTICE AND SIGNATURE PAGE

Using Government drawings, specifications, or other data included in this document for any purpose other than Government procurement does not in any way obligate the U.S. Government. The fact that the Government formulated or supplied the drawings, specifications, or other data does not license the holder or any other person or corporation; or convey any rights or permission to manufacture, use, or sell any patented invention that may relate to them.

This report was cleared for public release by the 88th Air Base Wing Public Affairs Office and is available to the general public, including foreign nationals. Copies may be obtained from the Defense Technical Information Center (DTIC) (<http://www.dtic.mil>).

AFRL-RH-WP-TR-2011-0019 HAS BEEN REVIEWED AND IS APPROVED FOR PUBLICATION IN ACCORDANCE WITH ASSIGNED DISTRIBUTION STATEMENT.

//SIGNED//
BONNIE D. RIEHL
Work Unit Manager
Sensemaking & Organizational
Effectiveness Branch

//SIGNED//
DAVID G. HAGSTROM
Anticipate & Influence Behavior Division
Human Effectiveness Directorate
711th Human Performance Wing
Air Force Research Laboratory

This report is published in the interest of scientific and technical information exchange, and its publication does not constitute the Government's approval or disapproval of its ideas or findings.

REPORT DOCUMENTATION PAGE				Form Approved OMB No. 0704-0188	
Public reporting burden for this collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing this collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden to Department of Defense, Washington Headquarters Services, Directorate for Information Operations and Reports (0704-0188), 1215 Jefferson Davis Highway, Suite 1204, Arlington, VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to any penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number. PLEASE DO NOT RETURN YOUR FORM TO THE ABOVE ADDRESS.					
1. REPORT DATE (DD-MM-YYYY) 05-01-2011		2. REPORT TYPE Final		3. DATES COVERED (From - To) August 2008 – January 2011	
4. TITLE AND SUBTITLE An Intelligent Decision Support System for Workforce Forecast				5a. CONTRACT NUMBER FA8650-08-C-6873	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER 62202F	
6. AUTHOR(S) Hung-da Wan, Fengshan F. Chen, Glenn W. Kuriger				5d. PROJECT NUMBER 7184	
				5e. TASK NUMBER D2	
				5f. WORK UNIT NUMBER 7184D214	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) The University of Texas at San Antonio One UTSA Circle San Antonio TX 78249				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING / MONITORING AGENCY NAME(S) AND ADDRESS(ES) Air Force Materiel Command Air Force Research Laboratory 711th Human Performance Wing Human Effectiveness Directorate Anticipate & Influence Behavior Division Sensemaking & Organizational Effectiveness Branch Wright-Patterson AFB OH 45433-7604				10. SPONSOR/MONITOR'S ACRONYM(S) 711 HPW/RHXS	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S) AFRL-RH-WP-TR-2011-0019	
12. DISTRIBUTION / AVAILABILITY STATEMENT Distribution A: Approved for public release; distributed is unlimited.					
13. SUPPLEMENTARY NOTES 88ABW/PA cleared on 23 Feb 2011, 88ABW-2011-0727.					
14. ABSTRACT The main objective of this project is to enhance the effectiveness of workforce forecast and deployment through innovative approaches of artificial intelligence. The research included a final document of conventional and innovative methods of workforce forecasting and a decision support software program incorporating advanced artificial intelligence techniques for workforce forecasting.					
15. SUBJECT TERMS Work Forecast, Artificial Intelligence, Genetic Algorithm, Neural Network					
16. SECURITY CLASSIFICATION OF: UNCLASSIFIED			17. LIMITATION OF ABSTRACT SAR	18. NUMBER OF PAGES 142	19a. NAME OF RESPONSIBLE PERSON Bonnie D. Riehl
a. REPORT U	b. ABSTRACT U	c. THIS PAGE U			19b. TELEPHONE NUMBER (include area code) NA

THIS PAGE LEFT INTENTIONALLY BLANK

TABLE OF CONTENTS

Section	Page
PART I – CONVENTIONAL AND INNOVATIVE METHODS OF WORKFORCE FORECASTING	1
1.0 SUMMARY.....	1
2.0 INTRODUCTION TO WORKFORCE PLANNING AND ANALYSIS	2
2.1 Workforce Planning.....	2
2.2 Workforce Analysis	3
2.2.1 Step 1: Analyze Demand.....	4
2.2.2 Step 2: Project Supply.	4
2.2.3 Step 3: Analyze Gap.....	5
2.2.4 Step 4: Develop Strategy.....	5
3.0 STATE-OF-THE-ART OF WORKFORCE ANALYSIS: A GENERIC OVERVIEW.....	7
4.0 DEMAND ANALYSIS	11
4.1 Quantitative Method for Forecasting.....	11
4.1.1 Trend Extrapolation.....	11
4.1.2 Data Mining – Neural Networks.	14
4.1.3 Data Mining – Decision Trees.	15
4.1.4 ARIMA.....	18
4.1.5 Causal Models.	22
4.1.6 Segmentation.	23
4.1.7 Rule Based Forecasting.	23
4.1.8 Simulation.	25
4.1.9 Nonlinear Forecasting Techniques.....	26
4.2 Qualitative Method for Forecasting.....	28
4.2.1 Delphi Technique.	28
4.2.2 Nominal Group Technique.....	30
4.2.3 Conjoint Analysis.....	32
4.2.4 Judgmental Bootstrapping.....	35
4.3 Scenario Specific Forecasting Technique(s) Selection Tree	36

5.0	SUPPLY ANALYSIS	38
5.1	External Environment	39
5.2	Internal Factors	40
5.3	Workforce Supply Analysis in Different Sectors	42
5.3.1	Public Sector.	42
5.3.2	Health-care Sector.	43
5.3.3	Miscellaneous.....	44
6.0	GAP ANALYSIS.....	46
7.0	CONCLUSION.....	49
	REFERENCES	50
PART II – DECISION SUPPORT SOFTWARE PROGRAM USING ARTIFICIAL INTELLIGENCE TECHNIQUES FOR WORKFORCE FORECASTING		65
8.0	SUMMARY	65
9.0	INTRODUCTION TO WORKFORCE FORECASTING	66
9.1	Background.....	66
9.2	Problem Statement and Objectives	66
9.3	Detailed Overview	67
9.3.1	Step 1: Demand Analysis	67
9.3.2	Step 2: Supply Analysis	68
9.3.3	Step 3: Gap Analysis	69
9.3.4	Step 4: Strategy Development.....	70
9.4	Solutions Outline	71
10.0	THE FORECASTING DECISION SUPPORT SYSTEM (FDSS) FORMULATION.....	72
10.1	Data Collection	73
10.2	Parameter Identification.....	76
10.3	Fuzzy Logic Controller	76
11.0	SOFT COMPUTING AND ARTIFICIAL INTELLIGENCE	78
11.1	Artificial Neural Networks	78
11.1.1	Working of Neurons within a Human Brain	78
11.1.2	Simulating Neural Networks for Computation and Learning:	79

11.2	Genetic Algorithm	82
11.2.1	String (Chromosome) Representation.....	83
11.2.2	Selection.....	83
11.2.3	Crossover.....	84
11.2.4	Mutation	84
11.3	Algorithm of Self Guided Ants.....	85
11.3.1	The Initialization Criteria	86
11.3.2	The Self Guided State Transition Rule	86
11.3.3	The Pheromone Update Rules.....	87
11.4	Self Guided Ant-based Genetically-Optimized-Neural-Network (SGA GONN) Algorithm for DFSS	88
11.5	Clonal C-Fuzzy Decision Tree (C^2 FDT)	90
11.5.1	Fuzzy C-Means (FCM) Clustering.....	91
11.5.2	Clonal Algorithm.....	92
11.5.3	Clonal Algorithm Guided FCM Clustering.....	94
11.5.4	Development of C^2 FDT	95
11.5.5	Use of C^2 FDT in Forecasting Problems.....	98
12.0	SOFTWARE DEVELOPMENT	99
13.0	RESULTS AND DISCUSSION	110
13.1	Overview of the Simulated Case Study	110
13.2	Survey Questionnaire.....	110
13.3	The Fuzzy Logic Controller.....	112
13.4	Implementing SGA GONN on the Dataset.....	116
13.5	Limitations of this Research	123
14.0	CONCLUSIONS.....	124
	REFERENCES	126
	LIST OF ACRONYMS	129

LIST OF FIGURES

Figure	Page
Figure 1: Four Phases of Workforce Planning.....	2
Figure 2: Four Key Steps for Conducting Workforce Analysis	3
Figure 3: Demand Analysis Techniques	11
Figure 4: Example of a General Decision Tree	16
Figure 5: Structure of the Rule Base.....	25
Figure 6: Scenario Specific Forecasting Technique(s) Selection Tree.....	37
Figure 7: Workforce Supply Analysis Model.....	41
Figure 8: Four Key Steps for Conducting Workforce Analysis	67
Figure 9: Workforce Supply Analysis Model.....	69
Figure 10: Modeling Challenges in Workforce Forecasting.....	72
Figure 11: Building Blocks of FDSS	73
Figure 12: The Fuzzy Logic Controller	77
Figure 13: Structure of Neuron	79
Figure 14: Neural Network Structure (Tripathi et al., 2010).....	81
Figure 15: A Threshold Element as an Analog to the Neuron.....	82
Figure 16: Pseudo Code of Genetic Algorithm.....	83
Figure 17: PPX Example	84
Figure 18: SGA GONN with Weights t_i' and w_i	89
Figure 19: Schematic Representation of Proliferation and Maturation (Shukla et al., 2009).....	93
Figure 20: Antigen-binding Sites (Shukla et al., 2009)	93
Figure 21: Flow Chart of Clonal Algorithm (Shukla, 2009)	94
Figure 22: Growth of Decision Tree by Expanding Nodes (Shukla and Tiwari, 2009)	95
Figure 23: Traversing a Decision Tree	98
Figure 24: MATLAB Standard IDE for Software Development.....	99
Figure 25: MATLAB Function “AFRLv10Fun.m” Structure	100
Figure 26: MATLAB Code for Selecting the Data File	100
Figure 27: Function “SimulateAll.m” Implementation	101
Figure 28: Function “SimulateOne.m” Implementation.....	102
Figure 29: Deployment Tool.....	103
Figure 30: Solution Explorer for the Software	104
Figure 31: “Mainform.cs” Design Implementation	105
Figure 32: Coding of “Mainform.cs”	106
Figure 33: “Questionnaire.cs” Form Implementation.....	107
Figure 34: Coding of the “Questionnaire.cs”	108
Figure 35: Publish Application	109
Figure 36: The Survey Form.....	111
Figure 37: Fuzzy Logic Controller Setup (MATLAB Fuzzy Logic Toolbox)	113
Figure 38: Defuzzified Output Corresponding to the Answer and Importance Value	114
Figure 39: The Fuzzy Rule Base.....	114
Figure 40: Customizing the GA Parameters	116
Figure 41: Customizing NN Parameters	117
Figure 42: Structure of C ² FDT for Workforce Category 1.....	118

Figure 43: Forecasting for Employee Category 1 Training Data Set by C ² FDT	119
Figure 44: Average Percent Error vs. Size of Data Set in Training of C ² FDT for Employee Category 1	120
Figure 45: Comparative Analysis of the Convergence Trends for Three Algorithms.....	120
Figure 46: SGA GONN for the Training Data Set (Workforce Category 1).....	121
Figure 47: SGA GONN for the Training Data Set (Workforce Category 2).....	121
Figure 48: Comparative Analysis for Three Algorithms over Workforce Category 1 (Testing Dataset)	122
Figure 49: Comparative Analysis for Three Algorithms over Workforce Category 2 (Testing Dataset)	122

LIST OF TABLES

Table	Page
Table 1: The Gap Analysis Process	46
Table 2: The Gap Analysis Form.....	46
Table 3: The Gap Analysis Process	70
Table 4: Quantitative Input Survey Form	74
Table 5: Parameter List.....	76
Table 6: Relationship Matrix between Parameters and Questions	112
Table 7: The Input Output Dataset Format	115
Table 8: Parameter Settings	117
Table 9: Description of C ² FDT for Category 1 Employee Data Set.....	119

ABSTRACT

This project studies workforce forecasting in two main aspects: (1) an extensive review of the existing methodologies and techniques and (2) an effort to develop a decision support system with models and software programming. The main objective of this project is to enhance the effectiveness of workforce forecasting and deployment through innovative approaches of artificial intelligence. The deliverables include a summary report of conventional and innovative methods of workforce forecasting (Part I) and a decision support software program using artificial intelligence techniques for workforce forecasting (Part II).

Part I, provided a thorough literature review of fundamental research and practices of demand and supply forecasting techniques for workforce analysis. Over 300 relevant literatures have been identified and reviewed by the project team, and 289 of them are covered in this report. Following is a summary of the main contents and contributions of Part I:

- This report provides an extensive review of the fundamental knowledge, applications, guidelines, and scope of use of various workforce forecasting methods.
- A decision tree has been proposed for selecting an appropriate forecasting technique based on available data and the needs.
- A list of 289 reference resources related to workforce forecasting and planning have reviewed, including books, academic periodicals, conference papers, white papers, technical reports, and web resources.

In Part II, a forecasting decision support system has been developed, which consists of a forecasting model powered by artificial intelligence, a data collection scheme that covers a full spectrum of conditions, a software program to realize the proposed model, and simulated cases. Following is a summary of the main contents and contributions of Part II:

- This report presents a more comprehensive set of questionnaire with 52 questions and 17 parameters to capture the human resource information within most organizations.
- This report presents an improved intelligent decision support system for workforce forecasting. After conducting several tests on various artificial intelligence techniques, a “*Self Guided Ant-based Genetically-Optimized-Neural-Network*” (SGA GONN) model is proposed, which is a robust solution methodology that outperforms the existing solution methodologies compared in this research.
- A software program has been developed that integrates MATLAB, MS Excel, and Visual C#. It has a graphical interface that allows customization of the algorithm and also helps in maintaining the database. Users with little knowledge the models will still be able to operate the software by simply following the instructions with suggested default settings.

THIS PAGE LEFT INTENTIONALLY BLANK

PART I – CONVENTIONAL AND INNOVATIVE METHODS OF WORKFORCE FORECASTING

1.0 SUMMARY

This project studies workforce forecasting in two main aspects: (1) an extensive review of the existing methodologies and techniques and (2) an effort to develop a decision support system with models and software programming. The main objective of this project is to enhance the effectiveness of workforce forecasting and deployment through innovative approaches of artificial intelligence. The deliverables include a summary report of conventional and innovative methods of workforce forecasting (Part I) and a decision support software program using artificial intelligence techniques for workforce forecasting (Part II).

Part I contains a thorough literature review of fundamental research and practices of demand and supply forecasting techniques for workforce analysis. Over 300 relevant literatures have been identified and reviewed by the project team, and 289 of them are covered in this report. The vast amount of literature demonstrates the importance and complexity of the research of workforce planning, especially workforce demand and supply forecasting. Many different techniques have been proposed to conduct the workforce forecasting, including quantitative algorithms and qualitative (or judgmental) decision making methods. Yet, every technique has its strength, weakness, and limitation. How do we choose the most appropriate forecasting method or a combination of methods for forecasting future workforce?

In order to address the challenges, this report reviews the state of the art of workforce planning and forecasting techniques and provides decision support information, such as how to use a model, when to use it, and the advantages and limitations of the model. Guidelines and examples of various workforce planning activities have been included to illustrate the use of the models and techniques as well as the way to use them. In addition, a scenario specific forecasting technique(s) selection tree has been proposed in this report to help decision makers select the desired models or techniques based on the availability and type of time-series and cross-sectional data. Finally, the sources of the original material can be found in the reference list.

Following is a summary of the main contents and contributions of Part I:

- This report provides an extensive review of the fundamental knowledge, applications, guidelines, and scope of use of various workforce forecasting methods.
- A decision tree has been proposed for selecting an appropriate forecasting technique based on available data and the needs.
- A list of 289 reference resources related to workforce forecasting and planning have reviewed, including books, academic periodicals, conference papers, white papers, technical reports, and web resources.

Part II presents a forecasting decision support system developed by the research team, which consists of a forecasting model powered by artificial intelligence, a data collection scheme that covers a full spectrum of conditions, a software program to realize the proposed model, and simulated cases.

2.0 INTRODUCTION TO WORKFORCE PLANNING AND ANALYSIS

Workforce planning is an organized process for identifying the number of employees, their mix and the types of skill sets required to accomplish the organization's strategic goals and objectives. Workforce analysis is a subset of workforce planning, which covers the demand, supply and gap analysis of workforce. This report focuses mostly on the demand and supply forecasting techniques for workforce planning.

2.1 Workforce Planning

Workforce planning is not just about the 3R's (i.e., Recruitment, Retention, and Retirement) but it starts with a well directed strategic plan, reliable and structured workforce data, a strong internal and external environmental scanning, and a keen awareness of trends. Following are few key factors affecting strategic workforce planning (Cotton, 2007):

- A more ethnically diverse workforce.
- Increasing age of skilled employees
- Workforce demographics
- Turnover rate and intention to quit
- Increased competition for highly skilled employees
- Workplace conditions and culture

By taking these factors into account, organizations should develop workforce plans ensuring that it has and will continue to have the right people with the right skills in the right job at the right time performing at their assignments efficiently and effectively. Workforce planning comprises with four phases (Keel, 2006). A pictorial form is presented in Figure 1 (Shukla, 2009).

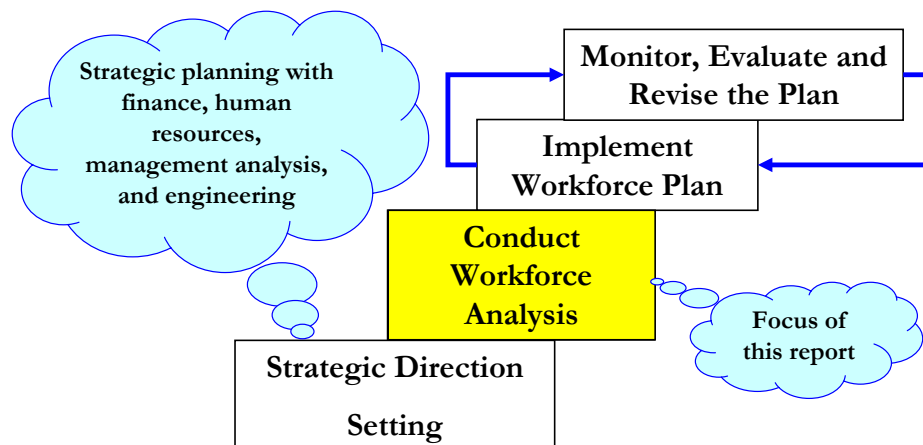


Figure 1: Four Phases of Workforce Planning

1. **Define the Organization's Strategic Direction:** Determine future functional requirements of the workforce through the organization's strategic planning and budgeting process.

2. **Conduct Workforce Analysis:** This is one of the most vital steps in effective workforce planning. It basically comprises the four following tasks:
 - (1) **Analyze demand:** Forecasting the size and mix of workforce composition, occupational competencies and proficiency levels needed to meet future mission requirements,
 - (2) **Analyze supply:** Determine current workforce competency and proficiency level profile and the projected profile in the future based on current trends and interventions,
 - (3) **Analyze gap:** A comparison of the Supply and Demand Forecasts to determine future shortages and surpluses in the workforce in terms of needed competencies,
 - (4) **Develop strategy:** Choose an appropriate set of accession, development, utilization, retention and separation strategies and timelines to resolve the most pressing gaps that will ensure an organization has a competent workforce to meet future mission needs.
3. **Implement Workforce Plan:** Communicate the workforce plan and Implement strategies, identified in phase 2 to reduce gaps and surpluses.
4. **Monitor, Evaluate and Revise:** After implementing the plan, assess what is working and what is not working and make adjustments to the plan accordingly. Finally, address new organizational issues that affect the workforce.

2.2 Workforce Analysis

Out of the abovementioned four phases, this literature review emphasizes step two (i.e., conduct workforce analysis). Analysis of workforce data is the key element in the workforce planning process. Workforce analysis frequently considers information such as occupations, skills, experience, retirement eligibility, diversity, turnover rates, education, and trend data. There are four key steps to the workforce analysis phase of the planning model. These steps are pictorially shown in Figure 2 (Keel, 2006).

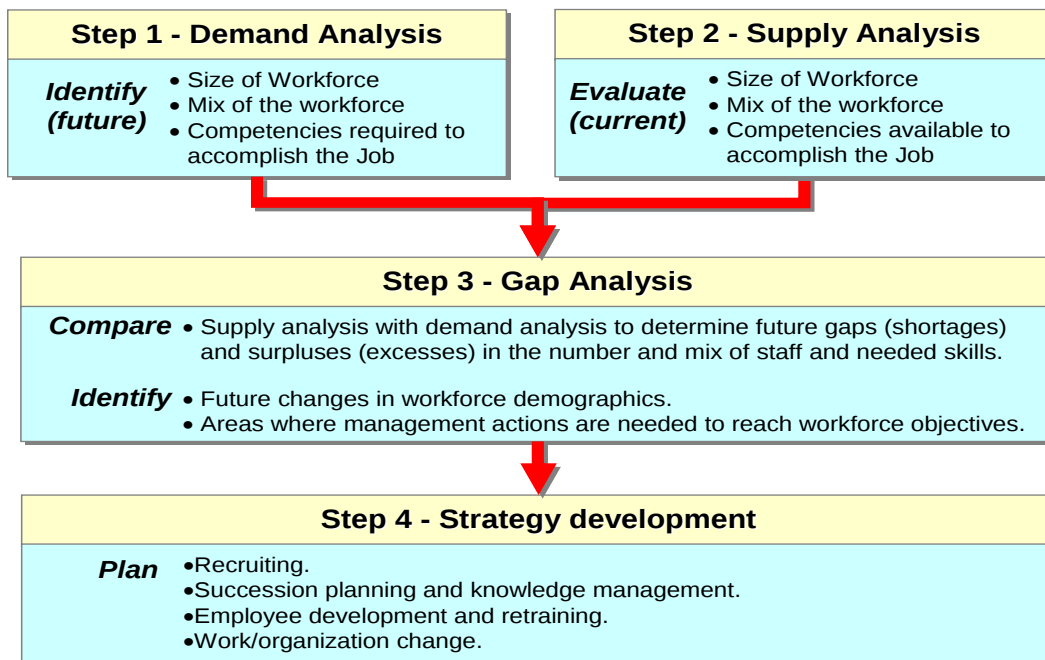


Figure 2: Four Key Steps for Conducting Workforce Analysis

2.2.1 Step 1: Analyze Demand

Forecasting the future workforce demand begins with the mission area analysis of the current products, processes, organization's policies and anticipating needed changes in the workforce to deliver future capabilities as a result of expected changes to those areas. The Demand Forecast will describe what each mission workforce (represented by positions) should be now (e.g., 1 to 5 years) in terms of:

- **Size:** total number of positions needed
- **Composition:** mix of workforce and contractor equivalents
- **Job requirements:** competencies required by crucial work aligned to core processes

In light of strategic direction and through collaboration with mission area experts, the true workforce requirement needed to accomplish future mission requirements (10 to 15 years) is assessed. Acknowledging the difficulty of describing the unknown, a disciplined balance of potential technological, threat, product, political, and economic changes with risk assessment is needed. The demand analysis process should examine not only what work the organization will do in the future, but how that work will be performed. Some possible considerations include:

- How will jobs and workload change as a result of technological advancements, economic, social, and political conditions?
- What are the consequences or results of these changes?
- What will be the reporting relationships?
- How will divisions, work groups, and jobs be designed?
- How will work flow into each part of the organization? What will be done with it? Where will the work flow?

Modeling the requirements determination process is the first step. Capturing the critical change factors is the second step and performing dynamic workforce simulations to facilitate decision-making and resource trade-offs is priceless.

Once the 'what and how' of future work has been determined through the identification of actual positions that are crucial to the core processes of the organization, the next step is to associate critical competencies that will be needed to carry out that work. The future workforce profile created through the Demand Forecast analysis will display a set of competencies that describes the future workforce. The Demand Forecast analysis should be as inclusive and transparent as possible. The workforce will have a greater understanding and ownership of the competency model if they are involved in the process. It will also give them a clearer idea of what the organization expects of a successful workforce. In addition, since developing the competency model is a visionary process, organizations should take care to include diverse viewpoints to avoid tunnel vision.

2.2.2 Step 2: Project Supply

The Supply Projection Model will be developed in the same manner. It will describe each mission workforce (the people) as it really is now and as it is likely to be in the future given existing personnel policies and practices. This task requires that the current and projected workforce be described in the same way that the ideal current and future workforce (positions) is

described on the demand side. It requires the application on an ageing model to determine the future workforce supply. Once the present workforce profile has been prepared, it will be projected out into the future as if no special management action was taken to replace attrition or develop existing workforce. Determining attrition rates for the organization and/or occupational areas and applying those to the present profile can accomplish this projection.

2.2.3 Step 3: Analyze Gap

Gap analysis is a process used to compare the workforce demand forecasted and the projected workforce supply. The resulting discrepancies are the identification of gaps and surpluses in personnel and competencies needed to carry out future mission objectives. Using the segmented approach, the health of each mission workforce is evaluated in terms of meaningful time horizons.

2.2.4 Step 4: Develop Strategy

The final step in the workforce analysis phase involves the development of strategies to address future gaps and surpluses. There is a wide range and combination of strategies that might be used to attract and develop the workforce with needed competencies, or to deal with excesses in competencies no longer needed for mission accomplishment. There is also a myriad of factors that will influence which strategy or, more likely, which combination of strategies that should be used. Some of these factors include, but are by no means limited to, the following:

- **Time.** Is there enough time to develop the workforce internally for anticipated vacancies or new competency needs, or is special, fast-paced recruitment the best approach?
- **Resources.** The availability of adequate resources will likely influence which strategies are used and to what degree, as well as priorities and timing.
- **Internal depth.** Does the existing workforce demonstrate the potential or interest to develop new competencies and assume new or modified positions, or is external recruitment needed?
- **"In-demand" competencies.** How high the competition is for the needed future competencies may influence whether recruitment versus internal development and succession is the most effective strategy, especially when compensation levels are limited.
- **Workplace and workforce dynamics.** Whether particular productivity and retention strategies need to be deployed will be influenced by workplace climate (e.g., employee satisfaction levels), workforce age, diversity, personal needs, etc.
- **Job classifications.** Do the presently used job classifications and position descriptions reflect the future functional requirements and competencies needed? Does the structure of the classification series have enough flexibility to recognize competency growth and employee succession in a timely fashion? Does it allow compensation flexibility?

Multiple options or solution sets for filling any gaps should be developed and priced. In this way, which gaps to resolve and which solutions to fund can be selected strategically – based on mission priorities, expected return on investment, and direct alignment to strategy.

In order to accomplish the above mentioned four steps, various methodologies and approaches exist in the literature. In subsequent sections we will discuss the methodologies available in the

literature for forecasting demand, accessing supply, measuring the gap and developing strategies for fulfilling this gap.

Regardless of the methods used, it is almost impossible to forecast the exact amount of future workforce in long planning horizons. There will always be an element of uncertainty until the forecast horizon has come to pass. Researcher and practitioners have put ample efforts in developing forecasting methods to minimize the element of uncertainty in forecasting. An intensive scanning of the literature shows that there is not much work regarding review of workforce forecasting methods. It is also seen that organizations feel the difficulty in selecting workforce forecasting methods according to their strategic direction and availability of data. There can be one or a combination of more than two suitable forecasting methods in a particular situation. Therefore, one of the biggest questions is how to choose the best forecasting method or combination of methods for forecasting the future workforce? For any organization, just to know the best suitable forecasting technique(s) is not enough for effective workforce forecasting. Identification of key factors affecting future workforce, questionnaires development for effective data collection, interpretation of responses to questionnaires, integration of responses to get final data are other vital tasks that should be performed in an organized way for effective and efficient workforce forecasting. In order to address above challenges in this report, first, the state-of-the-art of workforce planning/forecasting techniques is presented, and afterwards those are synthesized into a scenario specific forecasting technique(s) selection tree. Scenario specific technique(s) selection tree helps organizations in selecting the forecasting technique or combination of techniques suitable for workforce forecasting according to the availability and type of time-series and cross-sectional data.

3.0 STATE-OF-THE-ART OF WORKFORCE ANALYSIS: A GENERIC OVERVIEW

An intensive scanning of the literature reveals that there exists a plethora of research on workforce forecasting. These studies mainly focus on general labor market forecasting, and sector/occupation-specific forecasting. Unfortunately, in the literature, a large number of articles does not mention any specific models or approaches in their workforce demand analysis. Ward (1996) discusses the following six categories of workforce demand forecasting techniques: direct managerial input, best guess, historical ratios, process analysis, other statistical methods, and scenario analysis. Prescott (1991) focuses on the forecasting requirements for health care personnel. He developed analytic components for a comprehensive model by reviewing five general approaches and discussed their importance for understanding disequilibrium, particularly shortages, in the health care services providers. Chan et al., (2006) proposed a computer-based Auto-Regressive, Integrated, Moving-Average (ARIMA) model to forecast the demand for construction skills in Hong Kong. This model was based on workforce multiplier approach and used to predict the number of jobs created for a given level of investment. Moreover, they suggested that the model can be adapted by foreign construction authorities for manpower planning. Kovner and Reimers (1998) scrutinize data sources that can be used to forecast demand and supply for registered nurses in the United States. Anumba, et al., (2005) explored the potential of Geographic Information Systems (GIS) application to construction labor market planning. They suggested that GIS output needs to be in combination with a range of other forecasting techniques. A statistical measure which may be used to test the hypothesis that the forecasts provide an accurate picture of the structure of the labor market is proposed by Kolb and Stekler (1992). They used this measure to support the long-term forecasts of employment in different industries. Van der Laan (1996) reviews the five main classes of existing models which forecast regional supply and demand in the European Union. He also discussed the application and main features of the models, and an indication of the changes in economic direction of future model development.

Job Outlook Is Good (1999) presented a forecast of the expected growth of the number of industrial engineers in the USA by the year 2006. LeSage (1990) proposed an export-base error-correction model (ECM) in forecasting metropolitan employment. He provided an explanation of the time-series location-quotient derivation, Comparison of ECM with various types of vector autoregressive models, implications arise from the co-integration of the time series for export and local employment. Conway and Kniesner (1992) found that long-term contracting, adjustment costs, or efficiency wages can make a worker supply fewer (or more) hours than desired at the current economic opportunity set. A regression model with cross-correlated random coefficients is used to obtain new theoretically based measures of the direction and relative intensity of labor supply disequilibrium. Regressions incorporating person-specific fixed effects are consistent with a life-cycle labor supply model for salaried workers and suggest a need for fresh ways to include short-run labor demand forces in micro labor supply models of hourly paid workers.

Berkowitz (2004) reports the prediction of supply in pediatrics. He considered various conditions and situations that influence workforce projections, problems occurring in supply and demand for specialists and generalists and factors influencing the career choices of medical students. Li

and Dorfman (1995) developed an economic forecasting model to predict fluctuations in state-level employment growth. Brown (1999) developed a model to forecast dental workforce size and mix (by sex) for the first twenty years of the twenty first century in United States. According to his finding, there would be an increase in female dentists, a decrease in male dentists. Along with size and mix he also forecasted competencies required to deliver needed dental services. Labor market signaling approaches based workforce forecasting model was presented by Campbell (1997). He explained how labor market signal can be captured by monitoring workforce movements in the employment and unemployment of workers.

Cooper (1995) developed a model to forecast supply and demand of physicians in the United States for a period to 2020. Dumpe et al., (1998) recommends a forecasting model for the nursing workforce. Kolb and Stekler (1992) evaluated the long-term forecast of employment in various industries. They proposed a statistical measure which may be used to test the hypothesis that the forecasts provide an accurate picture of the structure of the labor market. The results were mixed with respect to the hypothesis, with the conclusions depending on the extent to which the data were rounded. A technique which could be used in evaluating long-term forecasts is proposed. This technique did not focus on the exact quantitative difference between the actual and predicted values. Rather, the distributions of the data were used. LeSage and Magura (1991) present the results of using input-output tables as a source of Bayesian prior information in a national employment forecasting model. A Bayesian vector autoregressive (BVAR) estimation technique is used to incorporate the inter-industry input-output relationships into the labor market forecasting model. This technique requires that a simple translation of the direct use coefficients from the input-output table be used as prior weighting elements to depict the inter-industry relations. The Bayesian model provides out-of-sample forecasts superior to those from unconstrained vector autoregressive, univariate autoregressive, a block recursive BVAR model and a naive BVAR model based on the Minnesota random walk prior.

Duffield and Coltrane (1992) test a model of the farm labor market for disequilibrium using CUSUM criteria. Borghans and Willems (1998), reveals that in manpower forecasting labor market developments are analyzed in terms of shortages and surpluses. Such an approach seems to neglect the flexibility of the labor market, present in the most economic labor market models. They have shown that an appropriate interpretation of gaps in manpower forecasting does not exclude a full functioning of the market clearing mechanism. Matarazzo (2000) offers a look at the demand and supply of library manpower from the 1970s to 1990s. He forecasted number of librarians that will reach 65 years of age by the year 2010 and vacancies from 1970 to 1985. Meehan and Ahmed (1990) put their effort to develop human resource forecasting model over the past decade by focusing on movement of people through the organization, and on large integrated models which are not always feasible or necessary in smaller organization. This paper presents the findings of a pilot study which utilized multiple regression demand models to forecast required staffing levels in an electrical utility company based on organizational variables, and to forecast the required mix of categories of employees.

Persad et al., (1995) present statistical models for forecasting the engineering manpower requirements for highway preconstruction activities. Two sets of models were developed in this study. In the first set of models, the independent variable used was an estimated construction cost, and in the second set of models the estimated construction cost and project type were the

independent variables. The models forecast engineering manpower requirements in terms of engineering man-hours as well as engineering cost. Rosenfeld and Warszawski (1993) present a systematic methodology for forecasting the demand for construction labor in various skills, within a national economy. Major factors, which determine the future needs for dwelling units and for other types of construction, are discussed in detail, while the strengths and weaknesses of different forecasting approaches are highlighted. Sweeney (2004) Detailed industry-occupation employment forecasts are an important class of regional labor market information produced by the U.S. Bureau of Labor Statistics. Eicher (1996) examines how interaction between endogenous human capital accumulation and technological change affects relative wages and economic growth. The proposed model provides a theoretical foundation for the empirically observed relation between technological change and relative demand, supply and wages of skilled labor. Tarlov (1995) focuses on Richard Cooper's article published in the November 15, 1995 issue of the "Journal of American Medical Association," on the projections in physician workforce in the United States by 2020. Importance of integrating into physician workforce; issue on the requirements for physician services.

Anumba et al., (2005) explores the potential of Geographic Information Systems (GIS) in providing such a mechanism for enhancing the labor market planning process. The paper details how GIS can aid construction labor market planning through its ability to integrate disparate labor market information efficiently, thereby placing analysts in a better position to understand specific spatial patterns.

Krolzig et al., (2002) presents a statistical model that offers a congruent representation of part of the UK labor market since the mid 1960s. They use a co-integrated vector autoregressive Markov switching model in which some parameters change according to the phase of the business cycle. The results of an impulse-response analysis highlight the dangers of using VARs when the constancy of the estimated coefficients has not been established, and demonstrate the advantages of generating regime dependent responses. Longhi et al., (2005) analyzed Artificial Neural Networks (ANNs) as a method to compute employment forecasts at a regional level. The empirical application is based on employment data collected for 327 West German regions over a period of fourteen years. First, the authors compare ANNs to models commonly used in panel data analysis. Second, they verify, in the case of panel data, whether the common practice of combining forecasts of the computed models is able to produce more reliable forecasts. Lutz (2002) pointed out that in western Germany the number of employed is likely to increase by 1.2 to 1.3 million between 2000 and 2015, for eastern Germany there are no indications of a positive labor market development with dynamics of its own. On the contrary, under "status-quo conditions" the number of employed in eastern Germany is likely to drop by 0.4 million in the period 2000-2015. This is the main finding of the latest IAB long-term projection on the basis of the IAB/INFORGE model. This model depicts the goods market for the whole of Germany by sectors in a highly disaggregated form. Popkov et al., (2005) developed a model based on description of cohorts competition and for labor demand—supply interaction. A modification of the random search is used for parametric identification of the model. Nonlinear dynamic systems with positive solutions and the parametric problem of entropy maximization (entropy operator) are considered. The continuity, differentiability and boundedness of the entropy operator and the boundedness of the solutions of the dynamic system are derived using the global implicit function theorem. The model is tested on real data from the EU countries.

Willems and de Grip (1993) revealed that replacement demand is an important component of the future demand for manpower often neglected in manpower demand forecasts. In particular, for characterizing the prospects for newcomers on the labor market, the manpower requirements method is not adequate as it merely focuses on employment mutations. This study builds a theoretical framework to measure the historical replacement demand distinguished by occupation and education. Roig (1999) studied the existence of two differentiated segments, primary and secondary, in the Spanish labor market. For this, a methodology (switching regression model with unknown regimes) is used which, on the one hand, does not demand a priori demarcation of segments and on the other, enables the allocation of workers to the segments to be treated as a factor endogenous to the model itself and closely connected with the wage-setting mechanisms operating in the segments. The empirical results show that the assignment to segments is not random and that there are substantial differences in the wage determination process between the two sectors.

4.0 DEMAND ANALYSIS

Forecasting the future workforce demand of an organization begins with the mission area analysis of its current products, processes, sub-organizations and policies and anticipating needed changes in the workforce to deliver future capabilities as a result of expected changes to those areas. Workforce demand forecasting mainly depends on type of organization hence there is not a uniform model to conduct it. In general, demand forecasting techniques can be broadly categorized into two types: judgmental (qualitative) and statistical (quantitative). The nature and complexity of the forecasting information available governs which one to be used from above two. Further, these two categories of forecasting techniques consist of various forecasting methods as shown in Figure 3 followed by more discussions.

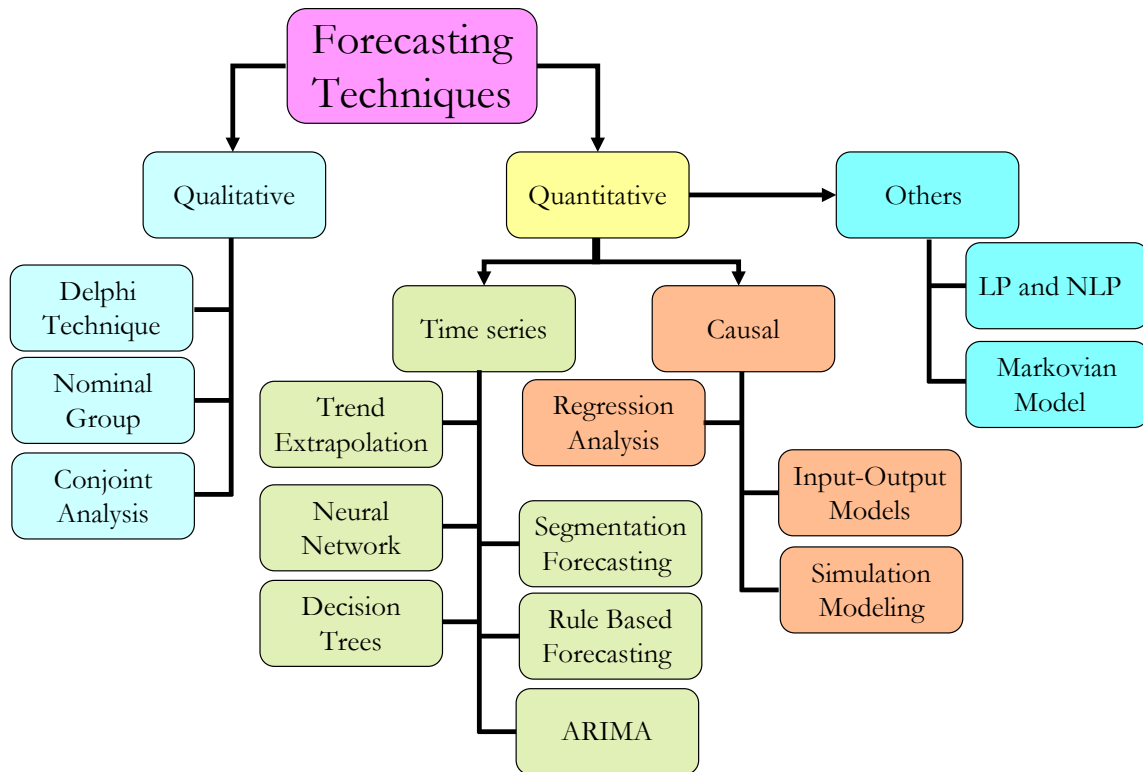


Figure 3: Demand Analysis Techniques

4.1 Quantitative Method for Forecasting

In this section, several methods based on quantitative data are reviewed. This is the main stream of forecasting techniques, and a large amount of models and techniques have been developed.

4.1.1 Trend Extrapolation

These methods examine trends and cycles in historical data, and then use mathematical techniques to extrapolate to the future. The assumption of all these techniques is that the forces responsible for creating the past will continue to operate in the future. This is often a valid assumption when forecasting short term horizons, but it falls short when creating medium and

long term forecasts. The further out we attempt to forecast, the less certain we become of the forecast. The extrapolation method is based on enlargement of studied developmental series (Makridakis et. al., 1998). It is based on an assumption that the studied process will develop in the future in the same direction or with the same intensity. Extrapolation is often shown graphically (as line graphs) with the level of a dependent variable on the y-axis and the time period on the x-axis. There are different "levels" of trends viz. constant, linear, exponential, and damped. Extrapolation consists of a number of principles a few are as follows (Armstrong, 2001a):

How to select and prepare data:

- Collect adequate amount of data pertaining to the situation that is to be forecasted. If sufficient data is not available then borrow data from analogous scenarios (Reilly and Chao, 1982). If similar situations cannot be found then conduct real time laboratory experimentation or simulation analysis.
- For the long term forecast use all relevant data (Smith and Sincich, 1990). There is lack of much evidence that a large amount of data guarantees an accurate forecast. While, in general, it is recommended to use all relevant data. Extrapolation implements the principle that recent data should be weighted more heavily and 'smoothes' out cyclical fluctuations to forecast the trend.
- With the divide-and-conquer strategy problems can be restructured to use domain expert's knowledge (MacGregor, 2001). To forecast workforce demand, break the problem down by part time, full time, civilians, military persons etc.
- Data preparation is a very important step to reduce measurement errors. Data preparation and cleaning is an often neglected but extremely important step in the extrapolation process. The old saying "garbage-in-garbage-out" is particularly applicable where large data sets collected via some automatic methods (e.g., via the Web) serve as the input into the analyses. Often, the method by which the data were gathered was not tightly controlled, and so the data may contain out-of-range values, impossible data combinations etc. Extrapolating the data that has not been carefully screened can produce highly misleading forecasts.
- Adjust data for historic events.

How to make extrapolation:

- Combine estimates of the level (Cairncross, 1969; Sanders and Ritzman, 2001)
- Use simple representation of trend unless there is strong evidence to the contrary (Makridakis and Hibon, 2000)
- Weight the most recent data more heavily than earlier data when errors are small, forecast Horizon is small and series is stable (Brown, 1962; Dalrymple and King, 1981)
- Use domain knowledge to provide pre-specified adjustments to extrapolation
- Use statistical methods as an aid in selecting and extrapolation procedure (Armstrong et. al., 2001)
- Update estimates of model parameters frequently (Tashman and Hoover, 2001)

Various methods of trend extrapolation:

- **Moving Average:** Extrapolation can be as simple as using demand in the current period to predict demand in the next planning horizon. For example, if demand is 100 units this

month, the forecast for next month's demand would be 100 units; if demand turned out to be 90 units instead, then the following month's demand would be 90 units, and so forth. This is sometimes referred to as naïve forecasting. However, this type of forecasting method does not take into account any type of historical demand behavior; it relies only on demand in the current period. As such, it reacts directly to the normal, random up-and-down movements in demand. Alternatively, the moving average uses several values during the recent past to develop a forecast. This tends to dampen, or smooth out, the random increases and decreases of a forecast that uses only one period. As such, the simple moving average is particularly useful for forecasting workforces that are relatively stable and do not display any pronounced behavior, such as a trend or seasonal pattern. The moving average method is good for stable demand with no pronounced behavioral patterns. Moving averages are computed for specific periods, such as 3 months or 5 months, depending on how much the forecaster desires to smooth the data. The longer the moving average period, the smoother it will be.

- **Weighted Moving Average:** The major disadvantage of the moving average method is that it does not react well to variations that occur for a reason, such as trends and seasonal effects (although this method does reflect trends to a moderate extent). Those factors that cause changes are generally ignored. It is basically a "mechanical" method, which reflects historical data in a consistent fashion. However, the moving average method does have the advantage of being easy to use, quick, and relatively inexpensive, although moving averages for a substantial number of periods for many different items can result in the accumulation and storage of a large amount of data. In general, this method can provide a good forecast for the short run, but an attempt should not be made to push the forecast too far into the distant future. The moving average method can be adjusted to reflect more closely more recent fluctuations in the data and seasonal effects. This adjusted method is referred to as a weighted moving average method. In this method, weights are assigned to the most recent data. In a weighted moving average, weights are assigned to the most recent data
- **Exponential Smoothing:** Smoothing and averaging are synonyms in forecasting; consequently, exponential smoothing might also be called exponential averaging. An exponentially smoothing value is actually a weighted moving average of all past actual values. Exponential smoothing refers to a set of methods of forecasting, several of which are still popular and widely used today.
 - o Single Exponential Smoothing (SES) is easy to apply because forecasts require only three pieces of data: (1) the most recent forecast, (2) the most recent actual observation, and (3) a smoothing constant. The smoothing constant determines the weight given to the most recent past observations and therefore controls the rate of smoothing or averaging. This method requires that the analyst choose a starting value to initialize the formula and an appropriate alpha (weight on actual observation).
 - o Brown's Double Exponential Smoothing uses a single coefficient, alpha, for both smoothing operations. This method computes the difference between single and double smoothed values as a measure of trend. It then adds this value to the

single smoothed value together with adjustment for the current trend. This method also requires starting values to initialize the formulas. When there is little historical data and the smoothing constant is small, choice of the initialization procedure can influence the fits and forecasts of several periods. When there are thirty or more periods of historical data, the starting values have little influence on period 31.

- o Holt's Two-Parameter Trend Model (double exponential smoothing model) uses a second smoothing constant, Beta, to separately smooth the trend. Holt's model further adjusts each smoothed value for the trend of the previous period before calculating the new smoothed value.
- o Winters' Three-Parameter Exponential Smoothing Model extends Holt's model by including a third smoothing operation and a third parameter to adjust for seasonality. The more parameters, the more decisions regarding weights and the initialization of parameters need to be made

When to use extrapolation:

- Many forecasts are required
- The forecaster is ignorant about the situation
- The situation is stable

4.1.2 Data Mining – Neural Networks

The world surrounding us generates various types of data in abundance. In the real life situation data arises from simulation, measurements, and centralization procedures. Most often, we meet with the paradox that more data means less information (Sorensen and Janssens, 2003). Recent developments in data storage devices, database management systems, computer technologies, and automatic learning techniques have made the data storage and its mining extremely economical and convenient. Data mining is a modern concept, generally employed to process the information embedded in the bulk of data. Recent developments in data mining techniques such as artificial neural networks, Bayesian networks, frequent patterns, decision trees, regression trees, and evolutionary algorithms have made the data mining as an independent new field of research. For a review on data mining techniques one may refer to (Chen et al., 1996).

Neural Networks are analytic techniques modeled according to the processes of learning in the cognitive system and the neurological functions of the brain. This is capable of predicting new observations (on specific variables) from other observations (on the same or other variables) after executing a process of so-called learning from existing data. Neural Networks is one of the data mining techniques. Rumelhart et al., (1986) discuss most of the neural network models in detail. Neural networks have been mathematically shown to be universal approximators of functions (Cybenko, 1989; Funahashi, 1989). This shows that neural networks can approximate whatever functional form best characterize the historical data. It might seem that because of their universal approximation properties neural networks should supersede the traditional forecasting techniques. That is not true for several reasons. One of the reasons is that universal approximation on a data set does not necessarily lead to good out of sample forecasts (Armstrong, 2001a). Hill et al., (1996) examines the literature on their forecasting performance. However, much data is needed to estimate neural network models and to reduce the risk of over-

fitting the data. But there is evidence that neural network models can produce forecasts that are more accurate than those from other methods (Adya and Collopy, 1998).

How to build neural network model for forecasting:

- Clean and scale the data prior to estimating the model (Kaastra and Boyd, 1996; Hill et al., 1996; Nelson et al., 1999)
- Use the right appropriate methods to choose the right straight point (Faraway and Chatfield, 1998; Marquez, 1992).
- Use specialized method to avoid local optima (Marquez, 1992; Sexton et al., 1998)
- Expend the network until there is no significant improvement in fit (Faraway and Chatfield, 1998; Zhang et al., 1998)
- Use pruning techniques when estimating neural networks and use the holdout samples when evaluating neural network (Marquez, 1992)
- Build plausible neural networks to gain model acceptance (Hill et al., 1996)
- Use following three approaches to ensure that the neural network model is valid (Adya and Collopy, 1998).
 - o Comparing the neural network forecast to the forecasts of other well accepted reference models
 - o Comparing the neural network and traditional forecasts' performance
 - o Marking enough forecasts to draw inference (say 40 forecasts)

When to use neural network models:

- Neural networks may be as accurate or more accurate than traditional forecasting methods for short-term (monthly and quarterly) forecasts (Foster et al., 1992; Hill et al., 1996)
- Neural networks may be better than traditional extrapolative forecasting methods for discontinuous series and often are as good as traditional forecasting methods in other situations (Armstrong and Collopy, 1992; Carbone and Makridakis, 1986)
- Neural networks are better than traditional extrapolative forecasting methods for long-term forecast horizons but are often no better than traditional forecasting methods for shorter forecast horizons (Sharda and Patil, 1992; Alekseev and Seixas, 2008)

4.1.3 Data Mining – Decision Trees

In recent years a number of classification techniques from both statistics and machine learning communities have been proposed (Fayyad et al., 1996; Quinlan, 1993). A well-known method of classification is the induction of decision trees (Breiman et al., 1984; Quinlan, 1993). Nodes of the trees are labeled with attribute, edges are labeled with possible values of corresponding attribute and leaves are labeled with different classes (Alexander and Grimshaw, 1996; Yildiz and Alpaydin, 2001). Objects are classified by following a path down the tree, taking the edges corresponding to the values of the attribute in an object. Decision trees (DTs) come with a comprehensive list of various training and pruning schemes, a diversity of discretization algorithms, and a series of learning refinements (Shukla and Tiwari, 2009; Breiman et al., 1984; Cantu-Paz and Kamath, 2000; Dobra and Gehrke, 2002; Pedrycz and Sosnowski, 2000; Weber, 1992). In order to classify an unlabeled data sample, the classifier tests the attribute values of the sample against the decision tree. A path is traced from the root to a leaf node, which holds the class predication for that sample. Decision trees can easily be converted into IF-THEN rules

(Quinlan, 1993) and used for decision-making. A generic structure of the tree is shown in Figure 4.

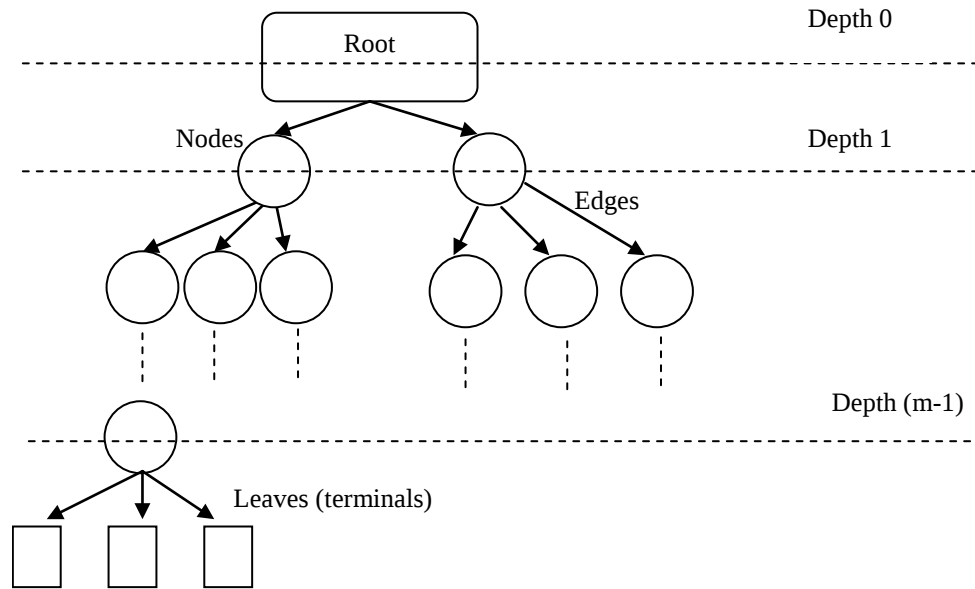


Figure 4: Example of a General Decision Tree

State-of-the-art of Decision Trees:

In the literature, various types of decision trees have been proposed according to the nature of the problem. One of the highly influential works on decision trees is the book written by Breiman et al., (1984), where a comprehensive treatment of decision tree techniques was presented and the classification and regression trees (CART) system was developed. Another significant contribution to decision tree techniques is Quinlan (1986; 1993), where ID3 (Iterative Dichotomiser 3) and C4.5 systems have been developed. Research interests in decision trees have been growing steadily since the mid 1980s. Most studies have focused on specific issues, such as splitting criteria, pruning methods, data representation, feature selection, etc. Various splitting criteria were proposed, including entropy or information gain (Hartmann et al., 1982; Quinlan, 1986), Gini index (Breiman et al., 1984), Twoing rule (Breiman et al., 1984), and its variant forms (White and Liu, 1994). The primary methods for pruning decision trees include cost-complexity pruning (Breiman et al., 1984), reduced-error pruning (Quinlan, 1987), and error-based pruning (Quinlan, 1993). In the literature DTs are classified roughly in following categories.

- **Univariate and Multivariate Decision Trees:** The ID3, C4.5, and CART decision trees and their variants are univariate trees. That is, at each node, a split of the data is made based on the value of a single variable. There have also been studies involving multivariate decision trees, which split the data based on the value of a linear combination of several variables. A multivariate decision tree system, named OC1, which combines deterministic hill-climbing with randomized procedures to search for a good split, has been developed by Murthy et al., (1994). A series of discriminant analysis based multivariate decision tree systems were developed by Kim and Loh (2001).

Another system based on linear discriminant analysis was developed by Gama and Brazdil (1999). Multivariate algorithms using linear machine decision trees were introduced in Draper et al., (1994). Linear-programming based multivariate tree algorithms were developed in Mangasarian (1997). More recent studies on multivariate decision trees can be found in (Setiono and Liu, 1999; Li et al., 2003; Lee, 2005).

- **Multiple Decision Trees:** One of the major limitations of decision tree is the high variance of the tree, specifically when data points are less and attribute values are more (Dietterich and Kong, 1995). In order to reduce the variance in classification performance Kwok and Carter (1990) and Buntine (1992) proposed the use of a pool of decision trees, instead of just one. Pool of the trees can be made by introducing randomness (Heath et al., 1993a) or by using different subsets of features for each individual tree (Shlien, 1992). Murphy and Pazzani (1994) proposed a decision forest consisting of decision trees built with a series of experiments. Moreover, they also investigated the functional relationship between the size of decision trees and accuracy of the tree on test data.
- **Incremental Decision Trees:** Fisher and Schlimmer (1988) developed ID4 (an incremental induction of decision trees) to avoid reconstruction of the decision tree, if in case, new training data is incorrectly or improperly classified by the prebuilt decision tree. Therefore, ID4 algorithm is able to build decision trees incrementally (i.e., existing decision tree can be updated for new instances). Utgoff (1989) suggests, ID5 or IDL, an advanced incremental algorithm that maintains statistics on the distributions of instances over attributes at each node in the tree in order to update the tree if necessary.
- **Fuzzy Decision Trees:** In the development of traditional decision trees the cut point is usually treated as crisp. However, due to the existence of vague and imprecise information in real-world problems, the class boundaries may not be defined clearly. In this case, the decision tree may produce high misclassification rates in testing even if they perform well in training (Cezary, 1998). This drawback could be overcome by firing multi-branches but with various certainty degrees. This could be implemented by means of fuzzy set theory, leading to fuzzy decision trees (Wang et al., 2000). Fuzzy decision trees simultaneously traverse multiple branches of a node with different degree of satisfaction ranged in $[0, 1]$. The induction of fuzzy decision trees follows the same steps as that of a classical decision tree with modified induction criteria (Janikow, 1998). Peng et al., (2000) proposed a criterion based on fuzzy mutual entropy in possibility domain. In these approaches, the continuous attributes are needed to be partitioned into several fuzzy sets prior to the tree induction, heuristically based on expert experiences and the data characteristics. Pedrycz and Sosnowski (2000; 2005) have developed cluster based fuzzy decision trees (CFDT). In the development of cluster based fuzzy decision trees fuzzy C-mean clustering were treated as its generic building blocks. Shukla and Tiwari (2009) developed genetically optimized cluster oriented soft decision trees (GCSDT) which was able to ameliorate the deficiencies allied with CFDT.

Moreover, apart from the developing variants of decision trees plenty of researches have been carried out to increase the efficacy and efficiency of Decision trees. Primary focuses of these researches are controlling tree size (Kim and Koehler, 1994; Auer et al., 1995; Kalkanis, 1993;

Esposito et al., 1997; Smyth et al, 1995), modifying test space (Fayyad and Irani, 1992; Utgoff and Brodley, 1990; Heath et al., 1993b; Zheng, 1995), modifying test search (Quinlan, 1986; 1987; Dietterich et al., 1996; Utgoff and Clouse, 1996; Pazzani et al., 1994), database restrictions (Wirth and Catlett, 1988; John, 1995) and alternative data structures (Oliver, 1993; Oliveira and Sangiovanni-Vincentelli, 1995; Kohavi and Li, 1995).

When to use decision trees:

- When instances can be described as attribute-value pairs
- Target function is discrete valued
- Possible noisy data
- When results are needed as a set of easily interpretable rules

4.1.4 ARIMA

Auto-regressive, integrated, moving-average (ARIMA) models are one of the most general paradigm for forecasting a time series. In fact, the easiest way to think of ARIMA models is as fine-tuned versions of random-walk and random-trend models. The fine-tuning consists of adding lags of the differenced series and/or lags of the forecast errors to the prediction equation, as needed to remove any last traces of autocorrelation from the forecast errors. The random walk model predicts the first difference of the series to be constant, the seasonal random walk model predicts the seasonal difference to be constant, and the seasonal random trend model predicts the first difference of the seasonal difference to be constant--usually zero. Lags of the differenced series appearing in the forecasting equation are called "auto-regressive" terms, lags of the forecast errors are called "moving average" terms, and a time series which needs to be differenced to be made stationary is said to be an "integrated" version of a stationary series. Random-walk and random-trend models, autoregressive models, and exponential smoothing models (i.e., exponential weighted moving averages) are all special cases of ARIMA models. AR, I, MA are the three phases of non-seasonal ARIMA model. This forecasting model works on following principles:

- Outcomes of the ARIMA depends on the linear functions of the sample observations
- The goal is to identify simplest paradigm that facilitates adequate description of the observed time series data.

ARIMA model is expressed as **ARIMA (p, d, q)**. In this representation:

- **p** is the number of autoregressive terms,
- **d** is the number of non-seasonal differences, and
- **q** is the number of lagged forecast errors in the prediction equation.
- **AR:** This component of the ARIMA model depicts how each observation is related to the previous p observations.
- **I:** This part of the ARIMA model determines whether the observed values are modeled directly, or whether the differences between consecutive observations are modeled instead. If $d = 0$, the observations are modeled directly.
- **MA:** This component of the ARIMA model depicts how each observation is related to the previous q errors (e.g., if $q = 1$, then each observation is a function of only one previous error).

To identify the appropriate ARIMA model for a time series, we begin by identifying the order(s) of differencing needed to stationarize the series and remove the gross features of seasonality, perhaps in conjunction with a variance-stabilizing transformation such as logging or deflating. If we stop at this point and predict that the differenced series is constant then we have merely fitted a random walk or random trend model. However, the best random walk or random trend model may still have auto-correlated errors, suggesting that additional factors of some kind are needed in the prediction equation.

ARIMA (0, 1, 0) (= random walk): Traditionally time-series models we encounter have two strategies for eliminating autocorrelation in forecast errors. One approach, which we first used in regression analysis, was the addition of lags of the stationarized series. For example, suppose we initially fit the random-walk-with-growth model to the time series Y. The prediction equation for this model can be written as:

$$\hat{Y}(t) - Y(t - 1) = \mu \quad (1)$$

where, the constant term (μ) is the average difference in Y. This can be considered as a degenerate regression model in which DIFF (Y) is the dependent variable and there are no independent variables other than the constant term. Since it includes (only) a non-seasonal difference and a constant term, it is classified as an "ARIMA (0, 1, 0) model with constant. Clearly, the random walk without growth would be just an ARIMA (0, 1, 0) model without constant.

ARIMA(1, 1, 0) (=differenced first-order autoregressive model): If the errors of the random walk model are autocorrelated, perhaps the problem can be fixed by adding one lag of the dependent variable to the prediction equation (i.e., by regressing DIFF(Y) on itself lagged by one period). This would yield the following prediction equation:

$$\hat{Y}(t) - Y(t - 1) = \mu + \phi (Y(t - 1) - Y(t - 2)) \quad (2)$$

which can be simplified as:

$$\hat{Y}(t) = \mu + Y(t - 1) + \phi (Y(t - 1) - Y(t - 2)) \quad (3)$$

This is a first-order autoregressive, or "AR (1)", model with one order of non-seasonal differencing and a constant term (i.e., an ARIMA (1, 1, 0) model with constant). Here, the constant term is denoted by " μ " and the autoregressive coefficient is denoted by " ϕ ", in keeping with the terminology for ARIMA models popularized by Box and Jenkins (1970).

ARIMA (0, 1, 1) without constant (= simple exponential smoothing): Another strategy for correcting autocorrelated errors in a random walk model is suggested by the simple exponential smoothing model. As we know for some nonstationary time series (e.g., one that exhibits noisy fluctuations around a slowly-varying mean), the random walk model does not perform as well as a moving average of past values. In other words, rather than taking the most recent observation as the forecast of the next observation, it is better to use an average of the last few observations in order to filter out the noise and more accurately estimate the local mean. The simple exponential smoothing model uses an exponentially weighted moving average of past values to achieve this effect. The prediction equation for the simple exponential smoothing model can be written in a number of mathematically equivalent ways, one of which is:

$$\hat{Y}(t) = Y(t-1) - \theta e(t-1) \quad (4)$$

where $e(t-1)$ denotes the error at period $t-1$. Note that this resembles the prediction equation for the ARIMA (1, 1, 0) model, except that instead of a multiple of the lagged difference it includes *a multiple of the lagged forecast error*. The coefficient of the lagged forecast error is denoted by the Greek letter " θ " (again following Box and Jenkins, 1970) and it is conventionally written with a negative sign for reasons of mathematical symmetry. " θ " in this equation corresponds to the quantity " $1-\alpha$ " in the exponential smoothing formulas.

When a lagged forecast error is included in the prediction equation as shown above, it is referred to as a "moving average" (MA) term. The simple exponential smoothing model is therefore a first-order moving average ("MA (1)") model with one order of non-seasonal differencing and no constant term (i.e., an "ARIMA (0, 1, 1) model without constant"). This means that statistical software that supports ARIMA models we can actually fit a simple exponential smoothing by specifying it as an ARIMA(0, 1, 1) model without constant, and the estimated MA(1) coefficient corresponds to " $1-\alpha$ " in the SES formula.

ARIMA (0, 1, 1) with constant (= simple exponential smoothing with growth): By implementing the SES model as an ARIMA model, we actually gain some flexibility. First of all, the estimated MA (1) coefficient is allowed to be negative: this corresponds to a smoothing factor larger than 1 in an SES model, which is usually not allowed by the SES model-fitting procedure. Second, you have the option of including a constant term in the ARIMA model if we wish, in order to estimate an average non-zero trend. The ARIMA (0, 1, 1) model with constant has the prediction equation:

$$\hat{Y}(t) = \mu + Y(t-1) - \theta e(t-1) \quad (5)$$

The one-period-ahead forecasts from this model are qualitatively similar to those of the SES model, except that the trajectory of the long-term forecasts is typically a sloping line rather than a horizontal line.

ARIMA (0, 2, 1) or (0, 2, 2) without constant (= linear exponential smoothing): Linear exponential smoothing models are ARIMA models which use two non-seasonal differences in conjunction with MA terms. The second difference of a series Y is not simply the difference between Y and itself lagged by two periods, but rather it is the *first difference of the first difference* (i.e., the change-in-the-change of Y at period t). Thus, the second difference of Y at period t is equal to $(Y(t)-Y(t-1)) - (Y(t-1)-Y(t-2)) = Y(t) - 2Y(t-1) + Y(t-2)$. A second difference of a discrete function is analogous to a second derivative of a continuous function: it measures the "acceleration" or "curvature" in the function at a given point in time.

The ARIMA (0, 2, 2) model without constant predicts that the second difference of the series equals a linear function of the last two forecast errors:

$$\hat{Y}(t) - 2Y(t-1) + Y(t-2) = -\theta_1 e(t-1) - \theta_2 e(t-2) \quad (6)$$

which can be rearranged as:

$$\hat{Y}(t) = 2Y(t-1) - Y(t-2) - \theta_1 e(t-1) - \theta_2 e(t-2) \quad (7)$$

where, “ θ_1 ” and “ θ_2 ” are the MA(1) and MA(2) coefficients. This is essentially the same as Brown's linear exponential smoothing model, with the MA (1) coefficient corresponding to the quantity $2(1-\alpha)$ in the LES model. To see this connection, we can check with forecasting equation for the LES model is:

$$\hat{Y}(t) = 2Y(t-1) - Y(t-2) - 2(1-\alpha)e(t-1) + (1-\alpha)^2 e(t-2) \quad (8)$$

Upon comparing terms, we see that the MA (1) coefficient corresponds to the quantity $2(1-\alpha)$ and the MA(2) coefficient corresponds to the quantity $-(1-\alpha)^2$. If α is larger than 0.7, the corresponding MA(2) term would be less than 0.09, which might not be significantly different from zero, in which case an ARIMA(0, 2, 1) model probably would be identified.

A "mixed" model – ARIMA (1, 1, 1): The features of autoregressive and moving average models can be "mixed" in the same model. For example, an ARIMA (1, 1, 1) model with constant would have the prediction equation:

$$\hat{Y}(t) = \mu + Y(t-1) + \phi(Y(t-1) - Y(t-2)) - \theta e(t-1) \quad (9)$$

Normally, though, we will try to stick to "unmixed" models with either only-AR or only-MA terms, because including both kinds of terms in the same model sometimes leads to overfitting of the data and non-uniqueness of the coefficients.

ARIMA models such as those described above are easy to implement on a spreadsheet. The prediction equation is simply a linear equation that refers to past values of original time series and past values of the errors. Thus, we can set up an ARIMA forecasting spreadsheet by storing the data in column A, the forecasting formula in column B, and the errors (data minus forecasts) in column C. The forecasting formula in a typical cell in column B would simply be a linear expression referring to values in preceding rows of columns A and C, multiplied by the appropriate AR or MA coefficients stored in cells elsewhere on the spreadsheet.

State-of-the-Art

Box and Jenkins (1970) developed a coherent generic three stage iterative cycle for time series identification, estimation and verification. With the evolution in computers technologies this iterative cycle (identification, estimation, and verification) was admired as an auto-regressive integrated moving-average (ARIMA) model. ARIMA model is also popularized as Box-Jenkins models. In the literature, ARIMA models have been studied extensively and treated as a major part of time series analysis. Gooijer and Hyndman (2006) states that, in general, ARIMA models can be divided in three categories viz. univariate ARIMA models, transfer function (dynamic regression) models, and multivariate (vector) ARIMA models.

The robustness of ARIMA models are reliant upon the ability of the model to mimic the behavior of diverse time series without requiring many parameters to be estimated in the final choice of model. Zellner (1971) found that the probability limit of the forecasts is largely influenced by parameter estimation. Box et al., (1994) describes number of methods for the estimation of ARIMA model's parameters. More recently, Kim (2003) considered parameter estimation and forecasting of AR (Auto-regressive) models in small samples. Hotta (1993) analyzed scrutinized the impact of an additive outlier on the intervals of forecasts when ARIMA model parameters are estimated. ARARMA methodology was proposed by Parzen (1982) as an alternative to ARIMA

models. The motivation behind developing the ARARMA was that a time series is transformed from a long memory AR filter to a short memory filter resulting in elimination of harsher differencing operators. Meade (2000) compared the forecasting performance of an automated and non automated ARARMA method.

The vector ARIMA (VARIMA), proposed by Quenouille (1957), is a multivariate generalization of univariate ARIMA model. Even VRIMA was proposed in 1975, but its software for implementation became available in early 90s. By considering that a stationary variable follows a VARIMA process Lutkepohl (1986) scrutinized the effects of temporal aggregation and systematic sampling on forecasting. Vector autoregressions (VARs) is a special case of VARIMA, which can be specified in a number of ways. Funke (1990) proposed five different VAR specifications and compared their forecasting performance by using monthly industrial production series. In general, VAR models tend to suffer from over-fitting with too many free insignificant parameters resulting in poor out of sample forecast. In order to overcome with the above mentioned difficulty Litterman (1986) imposed a prior distribution on the parameters with the view that many economic variables behave similar to random walk. Chevillon and Hendry (2005) studied the functional relationship of direct multi-step estimation of stationary and non-stationary VARs and forecast accuracy.

When to use ARIMA:

- When to forecast larger number of time series
- When much knowledge of the physical properties of the response variable are not available as ARIMA models are largely empirical or data driven
- When to do seasonal adjustment
- When to reduce revisions in the seasonally adjusted and trend estimates

4.1.5 Causal Models

Causal models are based on prior knowledge and theory. Time-series regression and cross-sectional regression are commonly used for estimating model parameters or coefficients. These models allow one to examine the effects of marketing activity, such as a change in price, as well as key aspects of the market, thus providing information for contingency planning. To develop causal models, one needs to select causal variables by using theory and prior knowledge. The key is to identify important variables, the direction of their effects, and any constraints. One should aim for a relatively simple model and use all available data to estimate it (Allen and Fildes, 2001). Surprisingly, sophisticated statistical procedures have not led to more accurate forecasts. In fact, crude estimates are often sufficient to provide accurate forecasts when using cross-sectional data (Dawes and Corrigan, 1974; Dana and Dawes, 2004).

Statisticians have developed sophisticated procedures for analyzing how well models fit historical data. Such procedures have, however, been of little value to forecasters. Measures of fit (such as R^2 or the standard error of the estimate of the model) have little relationship with forecast accuracy and they should therefore be avoided. Instead, holdout data should be used to assess the predictive validity of a model. This conclusion is based on findings from many studies with time-series data (Armstrong, 2001b). Statistical fit does relate to forecast accuracy for cross-sectional data, although the relationship is tenuous.

When to use causal models:

- Strong causal relationships are expected,
- The direction of the relationship is known,
- Causal relationships are known or they can be estimated
- Large changes are expected to occur in the causal variables over the forecast horizon, and
- Changes in the causal variables can be accurately forecast or controlled, especially with respect to their direction.

4.1.6 Segmentation

Segmentation involves breaking a problem down into independent parts, using data for each part to make a forecast, and then combining the parts. For example, a company could forecast sales of clothes separately for each climatic region, and then add the forecasts. To forecast using segmentation, one must first identify important causal variables that can be used to define the segments, and their priorities. For example, age and proximity to a beach are both likely to influence demand for surfboards, but the latter variable should have the higher priority; therefore, segment by proximity, then age. For each variable, cut-points are determined such that the stronger the relationship with dependent variable, the greater the non-linearity in the relationship, and the more data that are available the more cut-points should be used. Forecasts are made for the population of each segment and the behavior of the population within the segment using the best method or methods given the information available. Population and behavior forecasts are combined for each segment and the segment forecasts summed. Where there is interaction between variables, the effect of variables on demand are non-linear, and the effects of some variables can dominate others, segmentation has advantages over regression analysis (Armstrong, 1985).

When to use segmentation:

This is likely to occur where the segments are independent and are of roughly equal importance, and when information on each segment is good. Segmentation based on a priori selection of variables offers the possibility of improved accuracy at a low risk. Dangerfield and Morris (1992), for example, found that bottom-up forecasting, a simple application of segmentation, was more accurate than top-down forecasts for 74% of the 192 monthly time series tested. In some situations changes in segments are dependent on changes in other segments. For example, liberalization of gambling laws in city-A might result in decreased gambling revenue in already liberal cities B, C, and D. Efforts at dependent segmentation have gone under the names of microsimulation, world dynamics, and system dynamics. While the simulation approach seems reasonable, the models are complex and hence there are many opportunities for judgmental errors and biases.

4.1.7 Rule Based Forecasting

Rule-based forecasting (RBF) is a type of expert system that allows one to integrate managers' knowledge about the domain with time-series data in a structured and inexpensive way. For example, in many cases a useful guideline is that trends should be extrapolated only when they agree with managers' prior expectations. When the causal forces are contrary to the trend in the historical series, forecast errors tend to be large (Armstrong and Collopy, 1992). Although such problems occur only in a small percentage of cases, their effects are serious.

To apply RBF, one must first identify features of the series using statistical analysis, inspection, and domain knowledge (including causal forces). The rules are then used to adjust data, and to estimate short- and long-range models. RBF forecasts are a blend of the short- and long-range model forecasts.

RBF is most useful when substantive domain knowledge is available, patterns are discernable in the series, trends are strong, and forecasts are needed for long horizons. Under such conditions, errors for rule-based forecasts are substantially less than those for combined forecasts. In cases where the conditions were not met, forecast accuracy is not harmed.

Following are the steps for conducting rule base forecasting:

- ***Developing a Rule Base:*** Intention for rule-based forecasting is to apply a validated, fully disclosed, and understandable set of conditional actions to make forecasts. Although rule development begins by specifying rules that seem reasonable to experts, the goal is to go beyond such specifications and to validate rules by prior research and empirical tests. The rule base integrates several strategies for extrapolation. Among these are:
 - o using features of the series to establish weights for combining forecasts
 - o using heuristics to establish parameters for an exponential smoothing model
 - o using separate models for long-range and short-range forecasts
 - o damping the trend under certain conditions
 - o incorporating domain knowledge in extrapolation
- ***Attaining Information for Developing Rules:*** The forecasting literature provided some research streams that influence development of rules. Among these research streams were findings that favor decomposing time series, placing more weight on recent data, using simple methods, and making conservative trend estimates when uncertainty is high. Guidelines from the literature, however, typically lacked a precise statement of the conditions under which particular actions were likely to be useful. Also, such guidelines rarely specified precise actions needed to obtain forecasts.
- ***Specifying the Rules:*** Rules combine forecasts using weights that vary according to the features of the series. Given the desire for an understandable system, basically following simple, widely understood extrapolation methods are combined:
 - o The random walk emphasizes the short-range perspective; it sets the level to the last observation and is based on the assumption that there is no trend.
 - o The linear regression, which fits a least squares line to the historical data (or transformed historical data), represents the long range; we refer to its trend estimate as the basic trend.
 - o Holt's linear exponential smoothing captures information about short-range trends, and we call this the recent trend.
- ***Description of a Rule Base:*** One objective of the rule base is to provide more accurate forecasts. A second objective is to provide a systematic summary of knowledge. By expressing the knowledge explicitly and uniformly, we gain benefits generally associated with rule bases. These benefits include automating some tasks associated with

maintaining a complex body of knowledge and providing knowledge in an accessible and modifiable form. Besides its usefulness in forecasting, knowledge in this form aids reasoning about forecasting; that is, the knowledge is useful both procedurally and declaratively. Such explicit representation also can help in developing and testing theories. Figure 5 shows the elements of the rule base. First, features of the series are identified. Rules are then applied to produce short- and long-range forecasting models. To formulate these models we had to select smoothing factors for the exponential smoothing methods and make estimates of levels and trends for each model. For the long-range model, we formulated rules to damp the trend over the forecast horizon. Finally, there are rules for blending the forecasts - from the short- and long-range models.

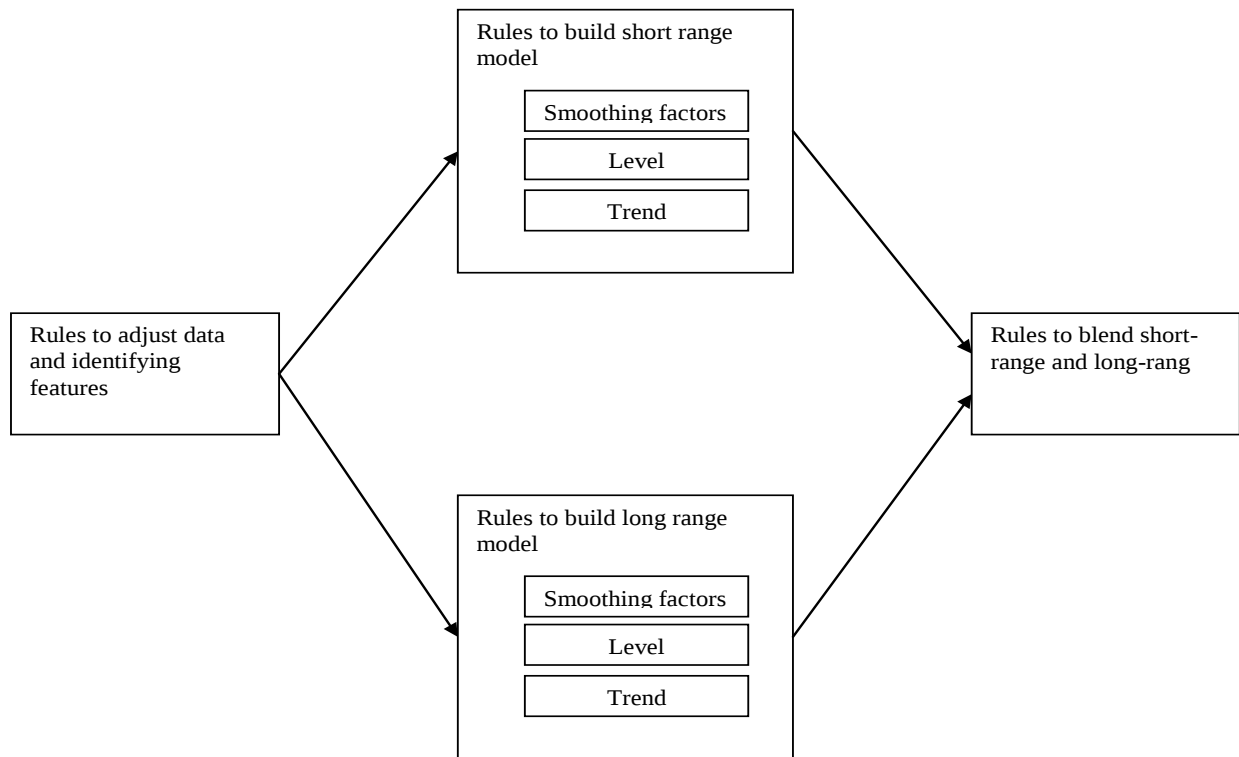


Figure 5: Structure of the Rule Base

4.1.8 Simulation

Computer simulation is the process of designing a model of real-life or theoretical physical system, executing it on a digital computer and analyzing the model output. Simulation model captures the mathematical and logical relationships among parameters, variables and other component parts of the system (Winston, 1994). Simulation models generally cannot be solved analytically like analytic models which can be computed by mathematical techniques viz. algebra, calculus or probability theory, (Law and Kelton, 1991). Simulation processes, same as conducting experiments on computers, deals with “how system would perform if ...” types of questions by displaying likely system’s performance under various input parameters settings. Simulation models are classified along three different categories (Law and Kelton, 1991).

- Simulations are called dynamic (or static) if they imitate system evolutions over time (or at a particular time point).
- Stochastic or deterministic depending on whether the model variables and parameters are probabilistic or known with certainty (Budnick et al., 1988).
- Discrete or continuous subject to whether model variables can change only at discrete time points or continuously over time (Law and Kelton, 1991).

Unlike analytical models, such as Markov chain models and optimization models, which can be represented by some fundamental equations, there are no universal mathematical or logical relationships to express simulation models since simulation by nature is system specific.

When to use simulation:

There are a number of reasons for using simulation versus analytical models which deal with the deficiencies of the models themselves such as:

- Lack of historical data
- Analytical models not available
- Existing analytical models are too complex
- Static results of analytical models are insufficient
- Analytical models only provide averages, not variability and extremes
- Analytical models cannot identify process bottlenecks or recommend design changes
- Analytical models often cannot provide sufficient detail nor identify interactions
- Animation is a better method of demonstrating results to management

Limitations:

While simulations are powerful and flexible in studying systems that are too complex for analytical models, the following limitations are noted:

- Constructions of simulation models, like design of experiments, can be time consuming and costly (Law and Kelton, 1991).
- For stochastic simulations, output data needs statistical analysis due to random variability (Winston, 1994). Output data from simulation are quite often auto-correlated, which requires special statistical techniques to make inferences. “There are still several output-analysis problems for which there is no completely accepted solution, and the methods available are often complicated to apply” (Budnick et al., 1988). Simulations are good at answering “what happens if ...?” type of questions but are not based on optimization algorithms (Law and Kelton, 1991).

4.1.9 Nonlinear Forecasting Techniques

The key concept in stochastic models is randomness, assuming that the process under study is governed by chance and probability laws. Based on a philosophy opposite to randomness, non-linear dynamic systems and chaos offer the possibility of describing complex phenomena as the result of a non-linear deterministic process. Compared to the study of linear time series, the development of nonlinear time series analysis and forecasting is still in its infancy. Although linearity is a useful assumption and a powerful tool in many areas, it became increasingly clear

in the late 1970s and early 1980s that linear models are insufficient in many real applications (Gooijer and Hyndman, 2006)

State-of-the-art of nonlinear forecasting techniques:

The idea of nonlinear time series forecasting has been proposed by Volterra (1930). He found that any continuous nonlinear function in t could be approximated by a finite Volterra series. Wiener (1958) extended ideas of functional series representation and further developed the existing material. Due to practical relevance of nonlinear forecasting, several useful nonlinear time series models were proposed in early 80s. De Gooijer and Kumar (1992) presented a review regarding the development of nonlinear forecasting techniques. At present, nonlinear forecasting is carried out by Monte Carlo simulation or by bootstrapping. The bootstrapping process is preferred since no assumptions are made about the distribution of the error process.

We have already discussed one of the widely used nonlinear forecasting models that is a neural network. Few other classes of nonlinear forecasting models are self-exciting threshold autoregressive (SETAR), continuous-time threshold autoregressive (CTAR), smooth transition autoregressive (STAR), functional coefficient autoregressive (FCAR) etc.

SETAR models were attributed by Tong (1983; 1990). Further, Clements and Smith (1997) compared a number of methods for obtaining multi-step-ahead forecasts for univariate discrete-time SETAR models. They concluded that forecasts obtained by using Monte Carlo simulation are satisfactory in cases where it is known that the disturbances in the SETAR model come from a symmetric distribution. Otherwise, the bootstrap method is to be preferred. Similar results were reported by De Gooijer and Vidiella-i-Anguera (2004). Brockwell and Hyndman (1992) obtained one-step-ahead forecasts for univariate continuous-time threshold autoregressive (CTAR). One drawback of the SETAR model is that the dynamics change discontinuously from one regime to the other. In contrast, a smooth transition autoregressive (STAR) model allows for a more gradual transition between the different regimes. Sarantis (2001) found evidence that STAR-type models can improve upon linear AR and random walk models in forecasting stock prices at both short-term and medium-term horizons. Functional coefficient autoregressive (FCAR) model is an AR model in which the AR coefficients are allowed to vary as a measurable smooth function of another variable, such as a lagged value of the time series itself or an exogenous variable. The FCAR model includes TAR and STAR models as special cases, and is analogous to the generalized additive model of Hastie and Tibshirani (1991). Chen and Tsay (1993) proposed a modeling procedure using ideas from both parametric and nonparametric statistics.

When to use nonlinear forecasting:

Davies et al., (1988) and Pemberton (1989), by numerical simulation, observe the phenomenon that the conditional mean and conditional median forecasts of nonlinear time series models have poor forecast performance compared to those of linear models. Therefore, if the linear forecast is not very much worse than the nonlinear forecast, then it is reasonable to adopt the former forecast rather than the nonlinear forecast, at least for computational reasons. More importantly, in practical applications for a given data set, unless an additional test is done in advance, there is no way of knowing whether the true model is factually linear or nonlinear. Thus, there arises the question that in what circumstances should the nonlinear predictors be adopted? By scanning the literature we found that in the following conditions non linear forecasting should be done:

- When historical data is not precise
- When time series is chaotic
- When system is nonlinear

4.2 Qualitative Method for Forecasting

In this section, several workforce forecasting methods based on qualitative judgment are reviewed. This category is very different from quantitative method. It addresses many of the qualitative aspects that are difficult for quantitative methods to handle.

4.2.1 Delphi Technique

The Delphi technique, mainly developed by Dalkey and Helmer (1963) at the Rand Corporation in the 1950s, is a widely used and accepted method for achieving convergence of opinion concerning real-world knowledge solicited from experts within certain topic areas. The Delphi technique is designed as a group communication process that aims at conducting detailed examinations and discussions of a specific issue for the purpose of goal setting, policy investigation, or predicting the occurrence of future events (Ulschak, 1983; Turoff and Hiltz, 1996; Ludwig, 1997). Common surveys try to identify “what is,” whereas the Delphi technique attempts to address “what could/should be” (Miller, 2006).

Theoretically, the Delphi process can be continuously iterated until consensus is determined to have been achieved. However, Cyphert and Gant (1971) and Custer et al., (1999) point out that three (or four) iterations are often sufficient to collect the needed information and to reach a consensus in most cases. In the first round, the Delphi process traditionally begins with an open-ended questionnaire. The open-ended questionnaire serves as the cornerstone of soliciting specific information about a content area from the Delphi subjects (Custer et al., 1999). In the second round, each Delphi participant receives a second questionnaire and is asked to review the items summarized by the investigators based on the information provided in the first round. Accordingly, Delphi panelists may be required to rate or rank-order items to establish preliminary priorities among items. In the third round, each Delphi panelist receives a questionnaire that includes the items and ratings summarized by the investigators in the previous round and are asked to revise his/her judgments (Jacobs, 1996). In the fourth and often final round, the list of remaining items, their ratings, minority opinions, and items achieving consensus are distributed to the panelists (Ludwig, 1994).

Conducting a Delphi study can be time-consuming. Specifically, when the instrument of a Delphi study consists of a large number of statements, subjects will need to dedicate large blocks of time to complete the questionnaires. Delbecq et al., (1975), Ulschak (1983), and Ludwig, (1994) recommend that a minimum of 45 days for the administration of a Delphi study is necessary. With regard to the time management between iterations, Delbecq et al., (1975) note that giving two weeks for Delphi subjects to respond to each round is encouraged.

Pioneering Delphi practitioners used varying panel size from few experts to several thousand peoples (Linstone, 1978). Linstone (1978) found that minimum panel size is seven and accuracy of forecast decorates very rapidly with the smaller size and improve slowly with large

size. Cavalli-Sforza and Ortolano (1984) concluded that typical size of Delphi panel should be eight to twelve members and Phillips (2000) found that optimum panel size should be from seven to twelve. While, Wild and Torgersen (2000) suggest panel size of 300-500 are usually considered useful.

State-of-the-art

The Delphi concept may be viewed as one of the spinoffs of defense research. "Project Delphi" was the name given to an Air Force-sponsored Rand Corporation study, starting in the early 1950's, concerning the use of expert opinion (Dalkey and Helmer, 1963). Starting in a nonprofit organization, Delphi has found its way into government, industry, and finally academe. This explosive rate of growth in utilization in recent years seems, on the surface, incompatible with the limited amount of controlled experimentation or academic investigation that has taken place. It is, however, responding to a demand for improved communications among larger and/or geographically dispersed groups which cannot be satisfied by other available techniques. Rowe and Wright (1999) systematically review empirical studies looking at the effectiveness of the Delphi technique.

In order to overcome the limitations of Delphi various modifications have been carried out. The Classical Delphi term is a method whereby, on an individual basis, data are collected from experts in a number of rounds (Rowe et al., 1991). At each stage, the results of preceding rounds are fed back until stability in responses among the experts on a specific issue has been reached through iteration. That is, no more significant changes occurring between rounds and Often consensus results (Zolingen and Klaassen, 2003). Gordon and Pease (2006) proposed real-time Delphi to replace the iterative process of the classical Delphi method. The advantages such as reducing costs and time significantly are obvious. A thorough comparison between classical round-based and real time Delphi is provided by Zipfinger (2007). Rauch (1979) developed the decision Delphi. He uses this type for decision making on social developments. The policy Delphi is widely used with social and political issues and is more suitable for application in the social sciences than the classical Delphi (Faché, 1993). The 'policy' Delphi also involves data being collected from experts on an individual basis in a number of rounds (Linstone, 1999). The traditional Delphi Method has always suffered from low convergence expert opinions, high execution cost, and the possibility that opinion organizers may filter out particular expert opinions. Murry et al., (1985) thus proposed the concept of integrating the traditional Delphi Method and the fuzzy theory to improve the vagueness and ambiguity of Delphi Method. Membership degree is used to establish the membership function of each participant. Ishikawa et al., (1993) further introduced the fuzzy theory into the Delphi term Method and developed max-min and Fuzzy Integration algorithms to predict the prevalence of computers in the future. However, this method is only applicable to the prediction of time series. Applicability of max-min Delphi method is analyzed with the forecast of possible timing for the realization of worldwide computerization forecast by specialists. The conclusion reached thereafter is applied to compare max-min Delphi method and conventional Delphi method. Besides, Hsu and Yang (2000) applied triangular fuzzy number to encompass expert opinions and establish the Fuzzy Delphi Method. The max and min values of expert opinions are taken as the two terminal points of triangular fuzzy numbers, and the geometric mean is taken as the membership degree of triangular fuzzy numbers to derive the statistically unbiased effect and avoid the impact of extreme values.

When to use Delphi Technique:

- The Delphi method is mainly used when long-term issues have to be assessed.
- It is suitable if there is the (political) attempt to involve many persons in processes.
- It tends to be used in evaluation when significant expertise exists on the subject.
- This method is recommended when the questions posed are simple (a program with few objectives, of a technical nature) and for the purpose of establishing a quantitative estimation of the potential impacts of an isolated intervention (e.g., increase in taxes or in the price of energy).
- It can be used to specify relations of causes and potential effects in the case of innovative interventions.
- It is particularly useful when a very large territory is being dealt with since there are no experts' travel expenses, only communication costs.

4.2.2 Nominal Group Technique

The Nominal Group Technique (see Delbecq et al., 1975, for a detailed review) was developed by Delbecq and Van de Ven in 1968 for conducting potentially problematic group meetings, embodies a set of procedures for conducting group sessions. Nominal Group Technique (NGT) generally consists of the following steps (Mahler, 1987; Moore, 1987; Delbecq et al., 1986):

- The moderator presents the question or problem to the group in written form and reads the question to the group.
- Participants independently and silently generate a list of ideas and write them down.
- Group members engage in a round-robin feedback session to concisely record each idea (without debate at this point). The moderator writes an idea from a group member on a flip chart that is visible to the entire group, and proceeds to ask for another idea from the next group member, and so on. Proceed until all members' ideas have been documented.
- Each recorded idea is then discussed to determine clarity and importance. For each idea, the moderator asks, "Are there any questions or comments group members would like to make about the item?" This step provides an opportunity for members to express their understanding of the logic and the relative importance of the item.
- Individuals vote privately to prioritize the ideas. The votes are tallied to identify the ideas that are rated highest by the group as a whole. The moderator establishes what criteria are used to prioritize the ideas.
- The moderator may shuffle the cards, redistribute, and ask participants to read results as s/he records them on the board.
- Optional discussion of results and revote

It is a popular technique and one of the most successful processes for structuring group meetings (Moore, 1990).

State-of-the-art

The Nominal Group Technique (NGT) has gained considerable recognition and has been widely applied in health, social service, education, industry, and government organizations. By imposing a precise structure upon an interacting group, NGT is intended to reduce process losses.

Gustafson et al., (1973) conducted some tests and results supported the potential accuracy of group judgments produced by a NGT approach. Subsequent studies have provided contradictory

evidence about the usefulness of NGT. The NGT presents several advantages to the researcher. It offers a well tested procedure which in one study was shown to be more effective at producing ideas than either the Delphi technique or discussion groups (Nemiroff et al., 1976).

Green (1975) conducted a study of NGT vs. normal interactive groups to determine problems faced by students in electronic data processing. He found that there was no significance between the groups as regards the number of ideas generated, the number of unique responses, or the quality of responses. Unique characteristics of Green's work were that he used a highly structured interactive group process to compare with nominal groups. Hegarty (1977) presents a summary of reasons for the superiority of nominal groups engaged in identifying problems, as regards three aspects of group problem identification activities.

The nominal group technique has been used in a wide variety of contexts related to different types of planning and control. These contexts range from the uncovering of productivity problems (Morris, 1979) to the identification of consumer perceptions of major pre-search problems (Claxton et al., 1980). Nemiroff et al., (1976) compared consensus, nominal and conventional interacting group on decision quality, member attitudes, and time taken to accomplish the lost at Sea exercise designed by the authors, but they noted no marked advantage to an NGT approach. Herbert and Yost (1979), strongly supported the superiority of NGT groups to uninstructed, interacting groups in accuracy, better use of group resources, and improvement over average individual decisions (high synergy) by implementing it on NASA Decision Making Problem (a task with considerable intentional depth). Lederer and Mendelow (1986) demonstrated the usefulness of the NGT in developing IS plans. Three nominal group technique sessions used IS practitioners from different levels of management to identify specific difficulties. By utilizing cards, Fox (1989) proposed the Improved Nominal Group Technique (INGT) that assures contributor anonymity, adds productive pre-meeting activity and removes NGT's inputting-transcribing bottleneck. He also concluded that INGT is appropriate for identifying and evaluating options, positions or problems, solving a problem, and for reviewing and refining written proposals or other documents.

Henrich and Greene (1991) used the NGT as one element of an action research program with a Fortune 100 company to facilitate a Manufacturing Resource Planning (MRP II) implementation. The results of a NGT with the project team were that top management involvement was seen as the most critical roadblock to the implementation. Based on the structure of NGT, Reisman et al., (1992) developed Group Decision Program (GDP), an interactive videodisc-based GDSS. Gallagher et al., (1993) explained the stages involved in conducting a nominal group and practical problems of its use in a health care setting are discussed with reference to a study of the priorities of care of diabetic patients, careers and health professionals. They also presented some potential applications of the technique in audit and exploratory research.

Dowling and Louis (2000) provide evidence that computer-assisted asynchronous (CAA) implementations of the NGT are more effective than non-computer-assisted synchronous (NCAS) implementations of the NGT (they generate more and better ideas, and do it in less time). In concrete, their implication is that organizations are wasting huge amounts of money on travel and accommodations for face-to-face meetings that could be conducted asynchronously. MacPhail (2001) describes the effective use of NGT with school-aged children to study factors

influencing curriculum choice. An exploratory study was conducted by Nelson et al., (2002) to use the NGT for determining the effectiveness and feasibility of 44 strategies for communication between home and school about homework assigned to students with high-incidence disabilities included in general education classrooms.

Tseng et al., (2006) studied an online NGT platform for implementing knowledge transfer. After doing experiments they concluded that an online NGT platform could provide formal activities for promoting knowledge transfer in pursuit of consensus at a distance. Lago et al., (2007) developed a methodology to evaluate the performance of NGT in a web-based environment compared to its traditional counterpart. Comparisons were made along several performance and process-related dimensions. The main conclusion was that, with regards to the decision process, participants felt that the traditional NGT outperformed the web-based NGT. With the help of NGT, Tuffrey-Wijne et al., (2007) presented a report of a study using the Nominal Group Technique as a method to elicit the views of people with intellectual disabilities on sensitive issues. They concluded that the NGT presents an effective and acceptable methodology in enabling people with intellectual disabilities to generate their views. Saremi et al., (2009) developed an approach for Total Quality Management consultant selection in Society of Manufacturing Engineers. The proposed method is based on TOPSIS (technique for order preference by similarity to ideal solution) method in fuzzy environment where decision criteria are obtained from the NGT.

When to use Nominal Group Technique:

- Group decision-making with difficult or non-conforming groups and under time pressure.
- Considering problems in staff meetings where political or status issues might reduce the input of some staff whose expertise is relevant.
- Obtaining solution ideas in a way that allows everyone to feel that they participated and were heard—people who want to dominate the discussion are reined in, people who would sit silently are brought out.
- Obtaining the input of a group while retaining the authority to make the decision independently.
- The group is less comfortable with more psychological methods.
- Individuals in the group may repress ideas because of timidity or dominance of others.

4.2.3 Conjoint Analysis

Conjoint analysis is used in marketing to study the factors that influence consumers' purchasing decisions (Green and Rao, 1971). It is used in designing new products, changing or repositioning existing products, evaluating the effects of price on purchase intent, and simulating market share (Louviere, 1988). Products possess attributes such as price, color, ingredients, guarantee, environmental impact, predicted reliability, and so on. Consumers typically do not have the option of buying the product that is best in every attribute, particularly when one of those attributes is price. Consumers are forced to make trade-offs as they decide which products to purchase. Researchers ask number of target market to indicate their choices for object under a range of hypothetical situations described in terms of product or service features, including attributes not included in the existing product or service. They use judgments to estimate preference functions, often a unique one for each respondent taking part in conjoint study. In essence, the researcher decomposes a respondent's overall preference judgments for objects

defined on two or more attributes for distinct attribute level. With the resulting performance functions researchers can forecast the share of preference for any product under consideration as compared to other products. By modifying the characteristics of a given product, the analyst can simulate a variety of plausible market situations. When used in this manner, conjoint analysis provides forecast of market share that allow managers to explore the market potential for new products or services. These forecasts, however, depend upon other factors such as availability of the products/services in a market scenario to each customer and the customers' awareness of these products/services. An application of conjoint analysis typically includes at least following steps (Wittink and Bergestuen, 2001):

- Selection of product/service category.
- Identification of a target market.
- Selection and definition of attributes.
- Selection of ranges of variation for the attributes.
- Description of plausible preference models and data collection methods.
- Development of the survey instruments.
- Sample size creation and data collection.
- Analysis of data.

State-of-the-art

From the early 70s, conjoint analysis has been considered as one of the major portfolios of techniques for measuring buyers' tradeoffs among multi-attributed products and services in both industry and academia (Green and Rao, 1971; Johnson, 1974; Srinivasan and Shocker, 1973). Green and Srinivasan (1978) described several advantages and limitations of the full-profile data collection method in comparison with the tradeoff procedure (two-factor-at-a-time method). Levy et al., (1983) described the telephone-mail-telephone procedure for collecting the data from the market. A variation on the -TMT theme is use of the locked box (Schwartz, 1978), a metal box containing all interview materials, which serves as a respondent premium. Steckel et al., (1990) review the literature on various ways of incorporating environmental inter attribute correlations in the construction of stimulus sets so as to increase the realism of the task. They provide a method for maximizing "orthogonality" subject to meeting various user-supplied constraints on the attribute levels that are allowed to appear together in full-profile descriptions. Krieger and Green (1988) and Wiley (1977) suggest methods for constructing stimulus sets for conjoint analysis that are Pareto optimal (i.e., no option dominates any other option on all attributes). Huber and Hansen (1986) and Green et al., (1988) report empirical results on the question of whether Pareto-optimally designed choice sets provide greater predictive validity than standard orthogonal designs in predicting a holdout set of realistic (Pareto-optimal) full profiles. Moore and Holbrook (1990) and Elrod et al., (1989) support and extend the findings of Green's et al., (1988) study. Johnson (1987) proposed computer based adaptive conjoint Analysis for paired comparisons of data collection.

Srinivasan et al., (1983) propose a constrained parameter estimation approach to improving predictive validity of conjoint analysis. Hagerty (1985) and Kamakura (1988) have proposed innovative approaches to improving the accuracy of full-profile conjoint analysis. Hagerty (1986) has argued that the emphasis on maximizing predictive power at the individual level may

be misplaced. He correctly points out that one should be more concerned with the accuracy of predicting market shares in the choice simulator.

Green (1984) states that the full-profile method of conjoint analysis works very well when there are only a few (say, six or fewer) attributes. Green and Srinivasan (1990b) reports that for handling large number of attributes three methods (1) the self explication approach, (2) hybrid conjoint analysis, and (3) Adaptive Conjoint Analysis (ACA) can be used. In the self-explication approach, the respondent first evaluates the levels of each attribute, including price, on a 0-10 (for example) desirability scale (with other attributes held constant) where the most preferred level on the attribute may be assigned the value 10 and the least preferred level assigned the value 0 (Green et al., 1981; Leigh et al., 1984; Wright and Kriewall, 1980). Hybrid models (Green et al., 1981) have been designed explicitly for task simplification in conjoint analysis. Hybrid uses self-explicated data to obtain a preliminary set of individualized part-worths for each respondent. The cross-validity of hybrid, traditional full-profile conjoint, and self-explication methods has been examined in three studies reported by Green (1984), three cases by Moore and Semenik (1988), and one by Akaah (1987). Adaptive Conjoint Analysis (ACA) from Sawtooth Software collects preference data in a computer-interactive mode. Johnson (1987) indicates that the respondent's interaction with the computer increases respondent interest and involvement with the task. The ACA system is unique in the sense that it collects (as well as analyzes) data by microcomputers. Two empirical studies have compared the ACA method and full-profile conjoint analysis (Agarwal, 1988). Green et al., (1990) report that the self-explication approach out predicted the ACA method in terms of cross-validated correlation and first choice hits, besides reducing the interview time considerably. For recent reviews of conjoint analysis please refer to Green et al., (2001).

Since the introduction of conjoint analysis in literature by Green and Rao (1971) as well as by Johnson (1974) it has been developed into a method of preference studies that receives much attention from both theoreticians and those who carry out field studies. Cattin and Wittink (1982) report 698 conjoint projects that were carried out by 17 companies in their survey of the period from 1971 to 1980. For the period from 1981 to 1985, Wittink and Cattin (1989) found 66 companies in the United States that were in charge of a total of 1062 conjoint projects. Wittink et al., (1994) counted a total of 956 projects in Europe carried out by 59 companies in the period from 1986 to 1991. Based on a 2004 Sawtooth Software customer survey, the leading company in Conjoint Software, between 5,000 and 8,000 conjoint analysis projects were conducted by Sawtooth Software users during 2003. The validation of the conjoint method can be measured not only by the companies today that utilize conjoint methods for decision-making, but also by the 171,000 hits on www.google.com. For more details of the statistics of conjoint analysis please refer to Gustafsson (2007). Refer to Green and Srinivasan (1990a) for detailed review.

Conjoint Analysis for the forecasting should be used when:

- Respondents represent a probability sample of the target market.
- Respondents are the decision makers for the product category under study.
- The conjoint exercise includes respondents to process information as they would in the marketplace.
- The alternatives can be meaningfully defined in terms of a modest number of attributes.

- Respondents find the conjoint task meaningful and have the motivation to provide valid judgments.

4.2.4 Judgmental Bootstrapping

Judgmental bootstrapping is a technique following the notion that decision makers usually have number of important factors in a task, but often they are not sure explicitly how much weight to give to each factor and are inconsistent in the weights they actually use (Tape et al., 1992). The appropriate weights can be extracted through regression on the past history of the task and used in a linear model to suggest good decisions or forecasts. In other words, Judgmental bootstrapping refers to developing a model of an expert's experience by regressing his forecasts against the information that he used. In bootstrapping, experts make predictions about real or simulated situations. A statistical procedure can then be used to infer the prediction model. It starts with the expert's forecasts and works backwards to infer rules the expert appeared to use in making these forecasts. It contrasts with the more common approach to experts systems, where one attempts to determine what rules were actually used and then perhaps what rules should be used. While, bootstrapping uses the expert's forecasts as the dependent variable and the information that the expert used serve as the causal variables. Following are few guidelines to develop judgmental bootstrapping:

- Consider all the process variables that the expert might use
- Quantify the process variables with a reasonable degree of accuracy
- Use the expertise of those forecasters who are most successful
- Ensure that the variables are valid
- Try to consider more than one expert or more than one groups of experts
- Use all expert's opinion if they differs in same situation
- Use large samples of stimulus cases
- Use stimulus cases that cover most reasonable possibilities
- Use stimulus cases that display low inter-correlations yet are realistic
- Use simple analysis to represent behavior
- Conduct formal monitoring

State-of-the-art

In the 1960s, researchers in a verity of fields studied judgmental bootstrapping. However, they were not aware of each other's work until Dawes (1971). Judgmental bootstrapping entered the managerial literature through the work of Bowman (1963) who argued that judgmental bootstrapping feedback led to better decisions primarily through the reduction of erratic decision making. He supported his argument with an analysis comparing the performance of actual decisions with the performance of decisions based on the bootstrapped model. There is a lot of empirical support for judgmental bootstrapping. Kleinmuntz (1990) provides an excellent review of this vast literature. Of particular interest is the work of Camerer (1981) who developed the conditions for judgmental bootstrapping models to outperform people; those conditions covering a wide variety of situations. One of the recent studies by Armstrong (2001) provides a meta-analysis of empirical comparisons on the value of judgmental bootstrapping as a means of improving judgmental forecast accuracy. (Ganzach et al., 2000) reports that judgmental bootstrapping has been found to be more accurate than unaided judgment in 8 of 11 comparisons, with two tests showing no difference, and one showing a small loss and the typical error

reduction was about 6%. Four of these bootstrapping studies were done in the last quarter century. They have helped to demonstrate the improved accuracy, extended the work to an applied management problem (e.g., advertising), and showed a condition under which it does not help. The failure occurred when experts used incorrect rules; in this case, bootstrapping applied incorrect rules more consistently and thus harmed accuracy Armstrong (2001). The principals for developing bootstrapping models are mainly based on the expert's opinion and on commonly accepted procedures in the social sciences and econometrics (Allen and Fildes, 2001). For more details of principles of model development one may refer to Armstrong (2001). Camerer (1981) discusses empirical evidence capturing a close difference between expert's predictions and those from their bootstrapping models. This does not imply that the bootstrapping forecasts are more accurate but it suggests that bootstrapping models capture key elements of an expert's decision process. Grove and Meehl (1996) did a remarkable review of the empirical evidences on econometric models and concluded that they are equal to or more accurate than unaided judgment in most settings. Ashton et al., (1994) find a bootstrap model out-performed by a statistical forecasting model. Lawrence et al., (2000) and Fildes et al., (2009) assert that bootstrapping is less effective in the context of time series extrapolation, where cue information tends to be autocorrelated; the "error bootstrapping" technique developed in response by Fildes (1991) seeks to model the errors in a judgmental forecast, much like the forecast correction approach described above.

Studies from various fields show bootstrapping improves upon accuracy of an expert's forecast. Judgmental bootstrapping models have been used in various fields like psychology, education, human resources, marketing, finance, and sales. Martorelli (1981) describes their use for draft selections in hockey and football. Christal (1968) reported that bootstrapping has been used for officer promotions in the U.S. Air Force. Batchelor and Kwan (2007) tested the proposition that decisions of experts are inferior to the decisions of statistical models of experts in the financial markets by using judgmental bootstrapping. O'Connor et al., (2005) conclude that integrating judgmental bootstrapping with task feedback support increased the forecasting accuracy and reliability. In essence, judgmental bootstrapping seems to be of benefit because judgmental inconsistency is eliminated (Newton, 1965; Stewart, 2001).

Judgmental Bootstrapping is used when:

- The problem is complex
- Reliable estimates can be obtained for bootstrapping
- Valid relationship are used
- The alternative is to use individual inexperienced experts

4.3 Scenario Specific Forecasting Technique(s) Selection Tree

Choosing the best forecasting method for any particular situation is a very tough task, and sometimes more than one method may be appropriate. In order to choose the best forecasting method among those available, this research proposes a decision making methodology, scenario specific forecasting technique (s) selection tree, shown in Figure 6. During the decision making via this tree, the first issue that needs to be resolved is what type of forecasting is required, qualitative, quantitative or both. If qualitative decisions are to be taken (e.g., competencies of workforce) then select left branch of the tree and traverse down. While, if quantitative decisions

(e.g., size and mix of the workforce) are to be made then select the right branch of the tree and traverse down. Moreover, there exist some situations where both type of decisions are required then call for both the edges at the first level of the tree. In case of quantitative procedures, the decision is further guided by whether the situation involves small or large changes. For small changes no policy analysis is needed and we obtain an acceptable feedback.

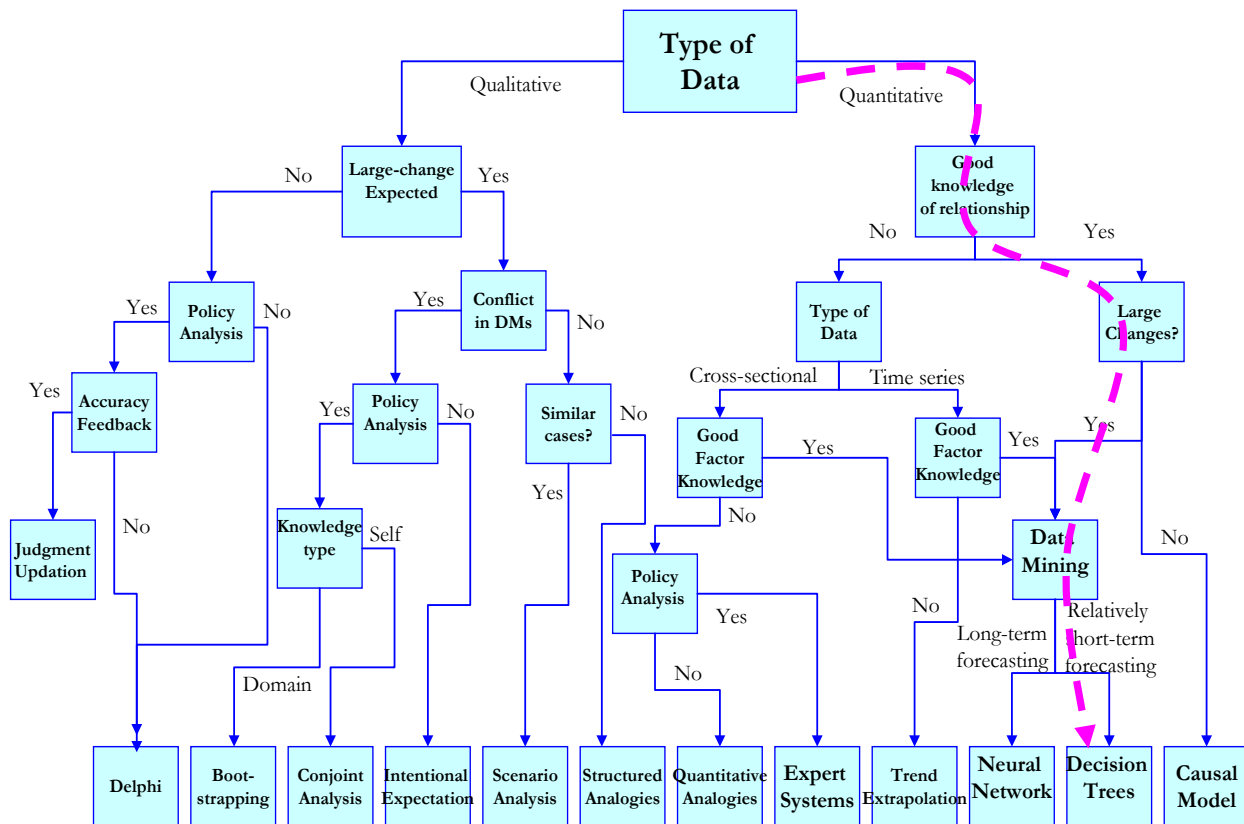


Figure 6: Scenario Specific Forecasting Technique(s) Selection Tree

If we have a lot of data, then it should be determined whether there is information about what empirical relationships might exist and if their magnitudes are known or not. If empirical knowledge of the relationship is available, then casual models are used. When large changes are improbable, we adopt data mining techniques. When, conditions are not clear, several appropriate techniques are combined to obtain the forecasts. Combined forecasts tend to achieve a high level of efficacy and reduce the likelihood of large errors. They are especially utilized when the component methods differ to a great extent among themselves. The integration of judgmental and statistical methods is summarized by Armstrong and Collopy (1998). Integration is efficient when judgments are collected in a systematically and then used as inputs to the quantitative models, rather than simply used as adjustments to the outputs.

5.0 SUPPLY ANALYSIS

Supply analysis accesses the available workforce size, composition, and competencies. In order to conduct this analysis organizations should consider workforce, workload, and competencies as integrated elements. The demographic data of the organization provides “snapshots” of the current workforce for the supply analysis process. Identifying employment trends is one of the major factors in projecting the future workforce supply. Organizations generally use transaction data to identify employment trends. Required transaction data can be collected by reviewing changes in workforce demographics by:

- Occupation
- Grade level,
- Organizational structure,
- Race/national origin,
- Gender,
- Age,
- Length of service and
- Retirement eligibility

This data also helps in developing valuable information on areas such as retirement eligibility or turnover for a given point in the future by projecting from current workforce demographic data. Personnel transaction data also provides the basis to identify baselines such as turnover rates. Moreover, it can also provide powerful tools to forecast workforce changes in the future that may occur from actions such as resignations and retirements. In conjunction with demographic data, transaction data help HR professionals and other managers to forecast opportunities for workforce change that can be incorporated into the action plans.

When modeling the current workforce, organizations must include permanent employees, supplemental direct hire employees, and contract workers (Cotton, 2007: Thompson and Mastracci, 2005):

- **Permanent employees** are on the organization’s payroll, have regular work hours, and are entitled to receive the benefits of regular employment offered by the organization.
- **Supplemental direct-hire employees** are on the payroll of the organization, but do not have regular work hours and are not entitled to full benefits of employment. Supplemental employees work on a temporary, seasonal, or on-call basis.
- **Contract workers are employees** whose “labor is procured through a contractual mechanism with a third party, such as a staffing agency.” They are employees who “work exclusively at the customer’s site, and whose work activities are integrated with those of the customer’s employees.

For examining the current supply organizations conduct SWOT (Strength, Weakness, Opportunity, and Threat) analysis. The SWOT analysis can be conducted by doing internal and external environmental scanning (IPMA –HR 2000). Environmental scanning leads to adaptive workforce plan in response to rapid workplace changes. Such scanning enables the managers to review and analyze internal and external Strengths, Weaknesses, Opportunities and Threats—the

SWOT analysis. Environmental scanning addresses external and internal factors that will affect short-term and long-term goals.

5.1 External Environment

Opportunities and threats created by key external forces that affect entire organizations should be examined. These key forces include demographics, economics, technology, and political/legal and social/cultural factors (relative to employees, customers and competitors). For example, environmental scanning will help in understanding recruitment/retention approaches and strategies that competitors currently use to attract hard-to-find specialists. Following are the external data that can be used for the SWOT analysis (IPMA –HR 2000):

- **General information such as:**
 - o Demand for and supply of workers in key occupational fields
 - o Emerging occupations and competencies
 - o Net migration patterns
 - o Retirement
 - o Desirability of key geographic areas
 - o Competitors in key geographic areas
 - o Policies of major competitors
 - o Labor force diversity
 - o Colleges' and educational institutions' enrollments and specialties
 - o New government laws and policies affecting the workforce
 - o General economic conditions
- **Changing composition of the workforce and shifting work patterns including demographics, diversity, outsourcing, work patterns, and work shifts such as:**
 - o Civilian labor force age
 - o Civilian labor force ethnicity
 - o Growing occupations/ethnicity in the civilian labor force
 - o Vanishing occupations/ethnicity in the civilian labor force
 - o Emerging competencies/ethnicity in the civilian labor force
 - o Civilian labor force education levels/ethnicity
 - o Civilian labor force secondary and post secondary school enrollments/ethnicity
 - o Civilian labor force high school graduations/ethnicity
 - o New social programs (e.g., school to work)
 - o Terminated social programs
 - o Current trends in staffing patterns (such as part-time or job sharing)
 - o Technology shifts
- **Government influences – policies, laws, and regulations affecting the workforce such as:**
 - o New employment laws
 - o Revisions in current employment laws
 - o Trends in lawsuits

- o Changes in rules and regulations (e.g., by the Environmental Protection Agency) that affect the work being studied
- **Economic conditions that affect available and qualified labor pools such as:**
 - o Unemployment rates (general)
 - o Unemployment rates in the specific geographic area of the organization
 - o Interest rates
 - o Inflation rates
 - o Interest rates in the specific geographic area of the organization
 - o Inflation rates in the specific geographic area of the organization
 - o Gross National Product trends
- **Geographic and competitive conditions such as:**
 - o Turnover data—general
 - o Turnover data—industry and occupation specific
 - o Secondary and post-secondary school enrollments
 - o Enrollments in curricula needed to support organizational strategies
 - o Net migration into the geographic area

5.2 Internal Factors

While it is important to identify threats and attractive opportunities in the external environment, it is even more critical to ensure that people and competencies are in place to meet those threats and take advantage of those opportunities. Identify internal strengths and weakness in light of the philosophy and culture of the agency. Internal data that agencies would look at to conduct a SWOT analysis are:

- Identify the current workforce skills, looking at education, language skills and competencies for successful performance.
- Identify retirement eligibility projections and patterns for key positions in the agency, specifically to determine where the agency is the most vulnerable to a wave of retirements and a loss of knowledge and the need for succession planning.
- Determine the demographic profiles of current employees, age, race, sex, etc., to determine the diversity of the workforce and areas for improvement.
- Determine the current state of the agency's union relations – is there a partnership relationship? If not, what will it take to develop a relationship that will support organizational change?
- Assess the organizational climate. Is the staff feeling supported and nurtured or are they feeling overwhelmed and burnt out or somewhere in-between? This assessment will help the agency understand where they need to begin in implementing change.
- Track turnover data to determine the amount of turnover in the agency, the types of turnover and reasons staff are leaving the agency to determine the impact turnover is having on the agency's ability to provide service.
- Understand the budget and the impact organizational change will have on salaries and benefits.

- Know the political environment. What might you expect in terms of possible changes in leadership: Governor, Commissioners, and Agency Director?

A SWOT analysis brings together the external and internal information to develop strategies and objectives. The SWOT analysis develops strategies that align organization strengths with external opportunities, identifies internal weaknesses, and acknowledges threats that could affect organization success. Of course, as with all analysis, budget considerations must be a major component of workforce supply analysis.

Workforce supply is, at its most basic level, the current workforce plus new hires less projected separations at some specific date in the future (Figure 7). For some organizations, projected workforce supply will be the result of a sophisticated mathematical model. For others, it will be an educated guess based upon data collected in the environmental scan. For most, it will be somewhere in between. No matter the level of sophistication, all models need to consider the same elements when projecting the future workforce composition: the inventory of the current workforce; the rate at which employees in specific occupations and at various leadership levels will leave the organization; what types of skills and abilities the organization will be able to attract; how the permanent workforce can be supplemented; and how the skill set of those who remain will or can change.

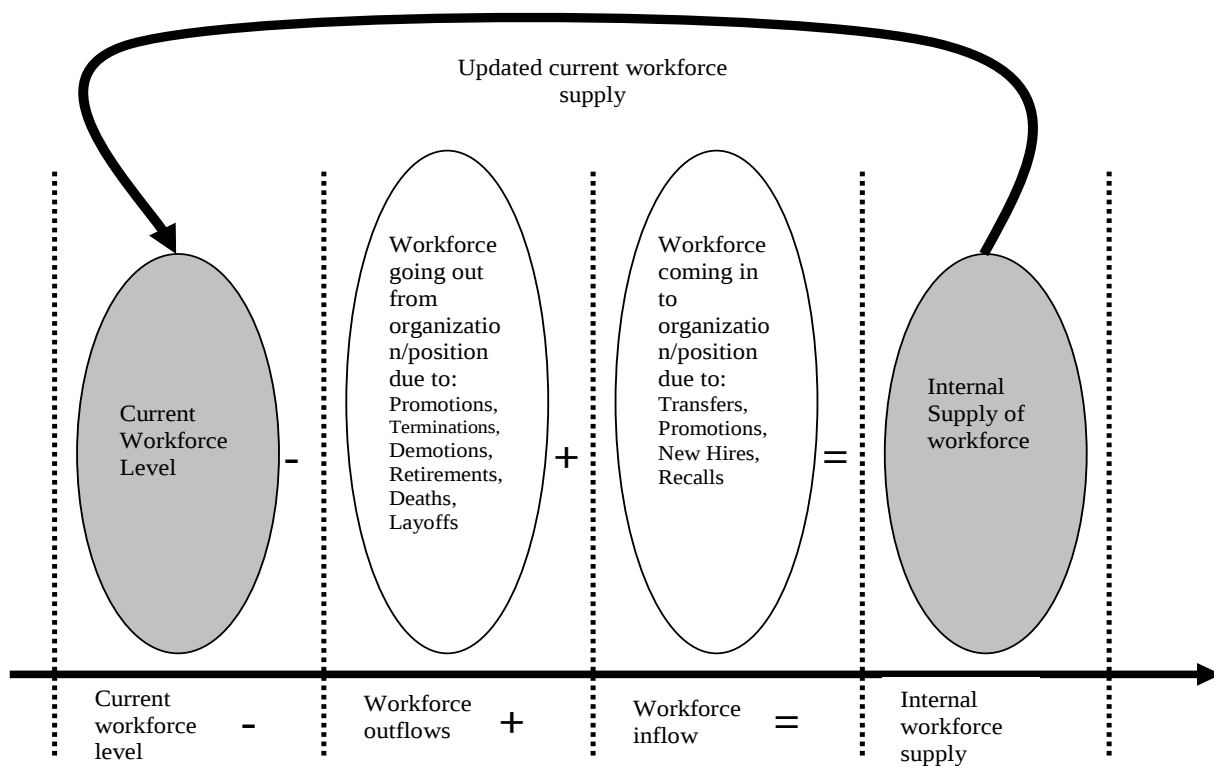


Figure 7: Workforce Supply Analysis Model

When computing how many people to add to the current workforce inventory, workforce planners must forecast who they will be able to recruit and what skills they will bring, how many

supplemental direct-hire employees will be available to the organization and what skills they bring, and what skills they will be able to acquire through the use of contract employees.

The forecast of new hires (the recruitment forecast) is arguably the most difficult component of the supply forecast. Whereas current and historical data provide a foundation for the current workforce inventory and the separation forecast, the recruitment forecast is influenced by many variables outside the control of the organization. From the environmental scan, the workforce planning team should be able to get a good idea of the general availability of the skill sets that will be needed. The most challenging part of the forecast is projecting the organization's ability to be a successful recruiter of the talent it needs. Recruitment success depends upon economic conditions, the local market for specific skills, and the competition for labor. While it may be hard to forecast precisely whom the organization will be able to hire, it should be relatively easy to identify positions that will be harder-than-average to fill. Another difficult component of the recruitment forecast is the unknown needed skill sets. No matter how well the organization plans, there is always a possibility that people will be needed to fill jobs that currently do not exist. The key in recruitment forecasting is to remember that it is a forecast—an imperfect projection of the future supply.

Projecting the availability of contract employees is another important component of the supply forecast. Because the companies for whom these employees work usually compete in the same labor market as the public sector organization, they can be subject to the same labor shortages and skill deficits as the public organization.

Once the current workforce inventory and recruitment forecast are complete, the next step is to forecast how many people will stay with the organization and why. Attrition stems from both voluntary separation (retirement, transfer, or resignation) and involuntary separation (termination for cause, layoff, long-term medical leave, death, or medical retirement).

Once again, for strategy development purposes, it is important to assess attrition at the micro and macro levels. In the case of voluntary separation, it is essential to collect information on why employees are leaving, where they are going, and what types of jobs they are taking. High turnover in a specific department or job classification may be a signal that salaries in the classification are too low, workload is too high relative to the salary, or working conditions are subpar. Micro-level knowledge is useful for developing targeted workforce planning strategies. The current workforce inventory must be compared to the assessment of the future workforce needs to identify the human capital gaps that must be filled to achieve organizational success.

5.3 Workforce Supply Analysis in Different Sectors

5.3.1 Public Sector

The present trend of workforce planning in government is based on ad-hoc and mixed strategies (Government Performance Project [GPP], 2001; International Personnel Management Association [IPMA], 2002; Selden et al., 2001). Although government units report having formal strategic planning in place, this is often not backed up with a formalized workforce plan (Kellough and Selden, 2003).

Choudhury (2007) states that in local government supply analysis involves in systematic mapping of the changing environment through monitoring demographic shifts, updating the recruitment process, and using attrition forecasts. He identifies, local officials view their environment as relatively stable, that small local governments find recruitment techniques only moderately important for securing supply. This indicates that environmental turbulence is yet to become a characteristic feature of small local governments. At the same time, this also points to the fact that local officials do not simply embrace the generic model of workforce planning but make realistic assessment of their need. In profiling their future, small local governments prefer specific techniques to anticipate future changes. In general, they consider analyzing the impact of technology and the state's budget on jobs to be more valuable than analyzing the impact of social and political trends. Their analysis includes not only the issue of competency but also the very nature of a jurisdiction's core functions.

Cotton (2007) practiced the workforce planning in U.S. Department of Transportation's (DOT's) and concluded that the federal human capital and workforce planning requirements is illustrative of the successes and challenges that occur when workforce planning is mandated and centrally supported. Cotton (2007) reported that DOT's success with workforce planning is evidenced by the fact that it was one of only five agencies to earn "green" status in the Human Capital Area. Report of the DOT (2006) states that, DOT performs an annual full-scope workforce analysis for workforce supply analysis as Human Capital Assessment and Accountability does. Cotton (2007) conducted the workforce supply analysis of U.S. DOT on a quarterly basis using Civilian Personnel Data File (CPDF) data and the Civilian Forecasting System (CIVFORS). Projections of workforce supply are also computed for all mission critical occupations. Mission-critical occupations are identified at the department and administration levels. Cotton (2007) also conducted the workforce supply analysis of the Maryland State Highway Administration in which, like many public sector organizations, lacked the comprehensive data sets needed to carry out supply analysis. However, there were able to produce a complete list of employees, their job titles, and assigned locations.

Keel (2006) suggests that supply analysis focuses on an agency's existing and future workforce supply. It answers the questions: what is the existing profile of the current workforce, and what does it need to be in the future to accomplish the agency's goals and objectives. Reviewing trend data will help an agency in projecting future workforce supply. It will also help an agency apply assumptions about how various factors will influence the future workforce. Trend information, combined with the current workforce profile, is an essential building block for forecasting workforce supply. To gather this information, some agencies have found it beneficial to delegate workforce planning to each division or satellite office. This gives managers the flexibility to address local issues, outcomes, and strategies.

Various reports on workforce supply analysis of the states (Arizona, Georgia, Iowa, Kansas, Washington, etc.) and local government exist in the literature. For more details one may refer to IPMA (2002).

5.3.2 Health-care Sector

In the health care sector workforce planning is one of the major crucial activities. In that sector, supply analysis model is termed as the trend model (Roberfroid et al., 2009). The trend model

relies on physician-population ratios and takes into account health care services currently delivered by the total pool of practicing physicians. This approach assumes that future requirements for physicians will need to match the volume of services currently provided on a per capita basis. O'Brien-Pallas et al., (2001) states that trend model reliant on three assumptions:

- the current level, mix, and distribution of providers in the population are adequate;
- the age and sex-specific productivity of providers remain constant in the future;
- the size and demographic profile of the providers change over time in ways projected by currently observed trends.

In such models, needs are defined as the necessary inflow of human resources to maintain or to reach at some identified future time, an arbitrary predefined level of service.

Although conceptually straightforward, such a model can gain complexity. First, the supply-based model often integrates parameters of demand. Possible changes in demographic features and the delivery system are sometimes factored into the projections. Second, the model is not necessarily based on a simple headcount of providers, but can integrate parameters linked to professional productivity.

The model can also serve to create scenarios, such as changes in the skill mix. In such instances, the model is called by some authors a substitution model (Persaud et al., 1999). The service targets approach is similar to the physician-to-population ratio. Requirements are defined on the basis of pre-set health service targets, e.g., staffing required for expansion of facilities (Zurn et al., 2004). The supply-based approach has been used in Belgium (Roberfroid et al., 2008), the United States of America (Lurie et al., 2002; Angus et al., 2000; Shipman et al., 2004; Holliman et al., 1997), Australia (Joyce et al., 2006), and Canada (Basu and Gupta, 2005).

Lavieri and Puterman (2009) used the attrition patterns to estimate the current supply of nurses for optimizing nursing human resource planning in British Columbia. Park and Choi (2001) estimated the supply of nurses in Korea by 2015 by using a baseline projection with an assumption that the current age structure of nurses is maintained. The age-specific future supply of nurses was estimated by using the age structure of input and output of nurses into the health care delivery system. Projections for each 5-year interval up to the year 2015 were made on the basis of the past trend of increments and losses before 1996. Maynard (2006) conducted the supply analysis by the activities of medical schools and open (that is, the EU in the case of Britain) and adjustable (external to the EU) migration trends.

5.3.3 Miscellaneous

Apart from public sectors, defense, healthcare, and industrial organizations workforce planning is widely practiced in fields like IT, banking, education, wild life etc. Limère and Landeghem (2008) developed a decision support system for workforce planning and analysis for a department of commercial credit applications of a large player in the Belgian banking sector. For conducting the supply analysis they scrutinized the existing electronic databases and performed the simulation analysis. National Wildlife Refuge System (2008) reported that Refuge System started conducting the first official phase of the seven step workforce planning in 2002. For the workforce supply analysis comprehensive demographic and skills information was collected in

Program Profiles for eight major program areas. This information served as a starting point for discussion and analysis at the Service's three-day Workforce Management Conference held in Washington, D.C. in May 2002. Refugee System leaders analyzed the information in the Program Profile and discussed the future demand/supply facing the Refugee System workforce. Ghosh (2005) analyzes the workforce supply in IT, health, consumer & commercial services, education, construction, engineering sectors of European Union (EU) states. He identified four major reasons (Population Trends, Participation rates in workforce, Participation rates in workforce, Education & Qualification) influencing workforce supply in EU labor markets.

Dubra and Gulbe (2008) identified and forecasted the Latvian labor market supply and demand of labor force in 15 sectors of national economy as well as determined causes of the disparity between the labor force supply and demand. The methodology of the studies is mainly based on data of different surveys and statistical information, expert evaluations, involving the use of methods of mathematical statistics and econometrics on a limited scale. More specifically, for conducting the supply analysis they analyzed following factors. (1) Analysis and forecasts for the demographic situation of Latvia, (2) Analysis of the annual CSB (Central Statistical Bureau) labor force survey results, (3) EU strategies in respect to labor market development and Euro-stat data, (4) Study proposal on the introduction of the Registers of Employees for the analysis and forecast of the labor market etc.

6.0 GAP ANALYSIS

Gap analysis involves in comparing the results of supply and demand analysis to identify gaps between the supply of and demand of total workforce and their mix. In gap analysis, managers use the workload and workforce data and the competency sets developed in earlier two phases (i.e., the supply and demand analysis phases). The most important thing to take care of in conducting gap analysis is the coordination in supply and demand data and competencies analyses as they have to be comparable.

Gap analysis identifies situations when demand exceeds supply such as when critical work demand, number of personnel, or current/future competencies will not meet future needs. It also identifies situations when future supply exceeds demand, however, such as when critical work demands, amount of personnel, or competencies exceed needs. In either event, your organization must identify these differences and make plans to address them. Future strategic planning direction will highly be determined from actions taken to eliminate the gaps. Depending upon how the supply and demand needs are determined and how specific they are, gaps can be identified by job title, series, grades, and locations. To be effective for comparison, the data and competencies in the supply and demand analysis phases need to be developed in tandem.

The “solution analysis” that will close the gaps must be strategic in nature. When doing solution analysis, organization should be prepared to address ongoing as well as unplanned changes in the workforce. The trends identified in supply and demand analysis can help your organization anticipate these changes.

In summary, calculating gaps will enable you to identify where your human capital (people and competencies) will not meet future needs (demand will exceed supply). The gap analysis also will determine whether your human capital exceeds the needs of the future (supply will exceed demand). There may also be situations where supply will meet future demand, thus resulting in a zero difference or no gap. The gap analysis process is outlined in Table 1 (IPMA –HR 2000).

Table 1: The Gap Analysis Process

How	What
Assess	The current supply of human capital
Factor in	Variables and assumptions
To come up with	Supply of human capital, then
Compare to	Demand
To come up with	Gaps and surpluses

Once you have identified the gaps between the demand (future needs and projected workload) and the supply (workforce and competencies), you will need to consult with management to set priorities to fill the gaps that will have a critical impact on your organization’s goals. Table 2 represents a form of conducting gap analysis.

Table 2: The Gap Analysis Form

WORKFORCE PLAN	Information Technology Specialist GS-2210-12 (Full Performance Level)
PRESENT SUPPLY	10
Minus Expected Attrition	-5
Plus Projected Supply (soon to be hired)	+3
FUTURE SUPPLY	8
Minus Future Demand	-12
WORKFORCE GAP	(4)

In the guide to workforce planning at the Department of Energy (2005), following questions are suggested to conduct the gap analysis:

- What will be the potential sources of your new staff that will be required?
- What attrition and retirement can be expected over next five years?
- Will attrition make it easier or harder to achieve workforce objectives?
- What kind of positions will need to be filled?
- How can training /re-training help?
- Succession planning implications?
- Competitive Sourcing solutions?
- Impact of budget decisions on any mission critical occupations?
- Any redeployment concerns or issues with current staff?
- Are new hires going to be required, and if so are they going to replace current employees or go into newly established positions?

Department of Navy (DON) (<http://www.doncio.navy.mil/workforce/>) Chief Information Officer (CIO) addressed requirement of gap analysis in its strategic goals, calling for the Navy and Marine Corps "to build IM/IT (Information Management/Information Technology) competencies to shape the workforce of the future." The CIO took specific action to achieve this goal by establishing a cross-functional, cross-organizational team - the IM/IT Workforce Integrated Process Team - chartered to define a workforce strategy and conduct IM/IT workforce planning analyses to facilitate more efficient and accurate alignment of the IM/IT workforce to meet the DON's organizational goals, commitments and priorities. In the report of CPS Human Resource Services (2007), "Gap Analysis: Staffing Assessment Template" is developed that synthesizes the data gathered during the Demand Analysis and Supply Analysis. The results from this analysis identify the gaps and surpluses in staffing levels needed for the future. Moreover, "Competency Gap Assessment Form" was also developed to identify gap of competencies.

Gates et al., (2006) highlights the fact that local DoD (Department of Defense) managers face a workforce-planning process that is substantially more complicated than the simple workforce-planning model would suggest. Managers must consider both the gap between workforce

demand and workforce supply and the gap between workforce supply and the workforce that can be supported with budgeted resources. If DoD wants managers to take requirements determination seriously, it must devise a way to eliminate the distinction between required and budgeted resources. It is possible that better DoD-wide data on workforce requirements could support this aim.

7.0 CONCLUSION

This project studies workforce forecasting in two aspects: (1) an extensive review of the existing methodologies and techniques and (2) an effort to develop a decision support system with models and software programming. In this report (i.e., Part I of this research) a thorough literature on fundamental research and practices of workforce planning has been presented. Workforce planning is an organized process for identifying the number of employees, their mix and the types of skill sets required to accomplish organization's strategic goals and objectives. The Part I research focuses mostly on the demand and supply forecasting techniques for workforce analysis, a subset of workforce planning. Over 300 relevant literatures (books, book chapters, journal papers, technical reports, etc.) have been identified and reviewed by the project team. 289 of the literatures are covered in this report as listed in the reference section.

The vast amount of literature demonstrates the importance and complexity of the research of workforce planning, especially workforce demand and supply forecasting. Many different techniques have been proposed to conduct the workforce forecasting, including quantitative algorithms and qualitative (or judgmental) decision making methods. Regardless of the methods used, it is almost impossible to forecast exact amount of future workforce in long planning horizons. There is always an element of uncertainty until the forecast horizon has come to pass. Therefore, one of the biggest questions is: How to choose the best forecasting method or combination of methods for forecasting future workforce? In order to address the challenges, this report reviews the state of the art of workforce planning/forecasting techniques and provides decision support information, such as how to use a model, when to use it, and the advantages and limitations of the model. Guidelines and examples of various workforce planning activities are included in the report to illustrate the use of the models and techniques as well as the way to use them. In addition, a scenario specific forecasting technique(s) selection tree has been proposed in this report to help decision makers select the desired models or techniques based on the availability and type of time-series and cross-sectional data. Finally, the sources of the original material can be found in the reference list.

In summary, the main contributions of this research include the following:

- This research provides an extensive review of the fundamental knowledge, applications, guidelines, and scope of use of various workforce forecasting methods. Each of the methods has been reviewed in terms of following questions:
 - o What is this method, and how does it work?
 - o How to use this method?
 - o When to use this method?
- A decision tree has been proposed for selecting an appropriate forecasting technique based on available data and the needs.
- A list of 289 reference resources related to workforce forecasting and planning have been included. The references consist of books, academic periodicals and conference papers, white papers, technical reports, and web resources. The sources of the reference are listed in the Reference section.

REFERENCES

- Adya, M. and Collopy, F. (1998). How effective are neural nets at forecasting and prediction? A review and evaluation, *Journal of Forecasting*, 17, 451-461.
- Agarwal, M. (1988). An empirical comparison of traditional conjoint and adaptive conjoint analysis (Paper No. 88-140). School of Management, State University of New York at Binghamton.
- Akaah, I.P. (1987). Predictive performance of hybrid conjoint models in a small-scale design: An empirical assessment. Wayne State University, Detroit, MI.
- Alekseev, K.P.G. and Seixas, J.M. (2008). A multivariate neural forecasting modeling for air transport – Preprocessed by decomposition: A Brazilian application, *Journal of Air Transport Management*, doi:10.1016/j.jairtraman.2008.08.008.
- Alexander, W.P. and Grimshaw, S. (1996). Treed regression. *Journal of Computational Graphical and Statistics*, 5, 156-175.
- Allen, P.G. and Fildes, R. (2001). Econometric forecasting. In J.S. Armstrong (Ed.), *Principles of forecasting: A handbook for researchers and practitioners* (pp. 303-362). Norwell, MA: Kluwer Academic Publishers.
- Angus, D., Kelley, M., Schmitz, R., White, A., and Popovich, J. (2000). Caring for the critically ill patient, current and projected workforce requirements for care of the critically ill and patients with pulmonary disease: can we meet the requirements of an aging population? *JAMA: Journal of the American Medical Association*, 284, 2762-2770.
- Anumba, C., Dainty, A.R.J., Ison, S.G., and Sergeant, A. (2005). The Application of GIS to construction labour market planning. *Construction Innovation*, 5(4), 219-230.
- Arizona. Salary plan and competing for talent, <http://www.hr.state.az.us/classcomp/ar2000.pdf>
- Armstrong, J.S. and Collopy, F. (1992). Error measures for generalizing about forecasting methods: Empirical comparisons, *International Journal of Forecasting*, 8, 69-80.
- Armstrong, J.S. (1985). *Long-term forecasting: From crystal ball to computer* (2nd ed.). New York: John Wiley.
- Armstrong, J.S. (2001). Judgmental bootstrapping: inferring expert rules for forecasting. In J. S. Armstrong (Ed.), *Principles of forecasting: A handbook for researchers and practitioners* (pp. 171-192). Norwood, MA: Kluwer Academic Publishers.
- Armstrong, J.S., (2001a). Evaluating forecasting methods. In J.S. Armstrong (Ed.), *Principles of forecasting: A handbook for researchers and practitioners* (pp. 365-382). Norwell, MA: Kluwer Academic Publishers.
- Armstrong, J.S. (2001b). Extrapolation of time-series and cross-sectional data. In J.S. Armstrong (Ed.), *Principles of forecasting: A handbook for researchers and practitioners* (pp. 217-243). Norwell, MA: Kluwer Academic Publishers.
- Armstrong, J.S. and Collopy, F. (1998) Integration of Statistical Methods and Judgment for Time Series Forecasting: Principles from Empirical Research. In G. Wright and P. Goodwin (Eds.), *Forecasting with judgment* (pp. 269-293). New York: John Wiley & Sons, Ltd.
- Armstrong, J.S., Adya, M., Collopy, F. (2001). Rule based forecasting: Using judgment in time-series extrapolation, In J.S. Armstrong (Ed.), *Principles of forecasting: A handbook for researchers and practitioners*. Norwell, MA: Kluwer Academic Publishers.
- Ashton, A.H., Ashton, R.H., and Davis, M.N. (1994). White-collar robotics: Levering managerial decision making. *California Management Review*, 37, 83-109.

- Auer, P., Holte, R.C., and Maass, W. (1995). Theory and application of agnostic PAC learning with small decision trees. In *Proceedings of the 12th International Conference on Machine Learning* (pp. 21-29). July 9-12, Tahoe City, CA. Also NeuroCOLT Technical Report NC-TR-96-034 (Feb 1996).
- Basu, K. and Gupta, A. (2005). A physician demand and supply forecast model for Nova Scotia, *Cahiers de Sociologie et de Demographie Medicales*, 45, 255-285.
- Batchelor, R. and Kwan, T.Y. (2007). Judgmental bootstrapping of technical traders in the bond market. *International Journal of Forecasting*, 23(3), 427-445.
- Berkowitz, C.D. (2004). Projecting, predicting, shaping: The challenge of workforce models. *Pediatrics*, 113(4), 918-919.
- Borghans, L. and Willems, E. (1998). Interpreting gaps in manpower forecasting models. *Labour*, 12(4), 633-641.
- Bowman, E.J. (1963). Consistency and optimality in managerial decision making. *Management Science*, 9(2), 310-322.
- Box, G.E.P. and Jenkins, G.M. (1970). *Time series analysis: Forecasting and control*. San Francisco: Holden-Day (revised ed. 1976).
- Box, G.E.P., Jenkins, G.M., and Reinsel, G.C. (1994). *Time series analysis: Forecasting and control* (3rd ed.). Englewood Cliffs, NJ: Prentice-Hall.
- Breiman, L., Friedman, J.H., Olshen, R.A., and Stone C.J. (1984). *Classification and regression trees*. Belmont, CA: Wadsworth.
- Brockwell, P.J. and Hyndman, R.J. (1992). On continuous-time threshold autoregression. *International Journal of Forecasting*, 8, 157-173.
- Brown, L.J. (1999). Trends in the dental health work force. *Journal of the American Dental Association*, 130(12), 1743.
- Brown, R.G. (1962). *Smoothing, forecasting and prediction*. Englewood Cliffs, NJ: Prentice-Hall.
- Budnick, F.S., McLeavey, D., and Mojena, R. (1988). *Principles of operations research for management*. Chicago, IL: Irwin.
- Buntine, W. (1992). Learning classification trees. *Statistics and Computing*, 2, 63-73.
- Cairncross, A. (1969). Economic forecasting. *Economic Journal*, 79, 797-812.
- Camerer, C. (1981). General conditions for the success of bootstrapping models. *Organizational Behavior and Human Decision Processes*, 27, 411-422.
- Campbell, C.P. (1997). Labour market signaling approaches. *Journal of European Industrial Training*, 21(8/9), 284-290.
- Cantu-Paz, E. and Kamath, C. (2000). Using evolutionary algorithm to induce oblique decision trees. In *Proceedings of the Genetic Evolutionary Computation Conference* (pp. 1053-1060). San Francisco, CA.
- Carbone, R. and Makridakis, S. (1986). Forecasting when pattern changes occur beyond the historical data. *Management Science*, 32, 257-271.
- Cattin, P. and Wittink, D. (1982). Commercial use of conjoint analysis: A survey. *Journal of Marketing*, 46(3), 44-53.
- Cavalli-Sforza, V. and Ortolano, L. (1984). Delphi forecasts of land use: Transportation interactions. *Journal of Transportation Engineering*, 110(3), 324-339.
- Cezary, Z.J. (1998). Fuzzy decision trees: issues and methods. *IEEE Transactions on Systems, Man, and Cybernetics-Part B*, 28(1), 1-14.

- Chan, A.P.C., Chiang, Y.H., Mak, S.W.K., Choy, L.H.T., and Wong, J.M.W. (2006). Forecasting the demand for construction skills in Hong Kong. *Construction Innovation*, 6(1), 3-19.
- Chen, M.S., Han, J., and Yu, P.S. (1996). Data mining: An overview from a database perspective. *IEEE Transactions on Knowledge and Data Engineering*, 8, 866-883.
- Chen, R. and Tsay, R.S. (1993). Functional-coefficient autoregressive models. *Journal of the American Statistical Association*, 88, 298-308.
- Chevillon, G. and Hendry, D.F. (2005). Non-parametric direct multistep estimation for forecasting economic processes. *International Journal of Forecasting*, 21, 201-218.
- Choudhury, E.H. (2007). Workforce planning in small local governments. *Review of Public Personnel Administration*, 27(3), 264-280.
- Christal, R.E. (1968). Selecting a harem and other applications of the policy-capturing model. *Journal of Experimental Education*, 36, 35-41.
- Claxton, J.D., Ritchie, J.R.B. and Zaichkowsky J. (1980). The Nominal Group Technique: Its potential for consumer research. *Journal of Consumer Research*, 7, 308-313.
- Clements, M.P. and Smith, J. (1997). The performance of alternative methods for SETAR models. *International Journal of Forecasting*, 13, 463-475.
- Collopy, F. and Armstrong, J.S. (1992). Rule-based forecasting: Development and validation of an expert systems approach to combining time series extrapolations. *Management Science*, 38, 1394-1414.
- Conway, K.S. and Kniesner, T.J. (1992). Estimating labour supply disequilibrium with fixed-effects random-coefficients regression. *Applied Economics*, 24(7), 781-789.
- Cooper, R.A. (1995). Perspectives on the physician workforce to the year 2020. *JAMA: Journal of the American Medical Association*, 274(19), 1534-1543.
- Cotton, A. (2007). *Seven steps of effective workforce planning*. Human Capital Management-Series (<http://www03.ibm.com/industries/government/us/detail/resource/X448444G98379A82.html>).
- CPS Human Resource Services (2007). *Workforce planning tool kit: Supply/Demand analysis and gap analysis*. Washington, DC: CPS Human Resource Services.
- Custer, R.L., Scarcella, J.A., and Stewart, B.R. (1999). The modified Delphi technique: A rotational modification. *Journal of Vocational and Technical Education*, 15(2), 1-10.
- Cybenko, G. (1989). Approximations by superpositions of sigmoidal functions. *Mathematics of Control, Signals, and Systems*, 2, 303-314.
- Cyphert, F.R. and Gant, W.L. (1971). The Delphi technique: A case study. *Phi Delta Kappan*, 52, 272-273.
- Dalkey, N. and Helmer, O. (1963). An experimental application of the Delphi method to the use of experts. *Management Science*, 9(3), 458-467.
- Dalrymple, D.J. and King B.E. (1981). Selecting parameters for short-term forecasting techniques. *Decision Sciences*, 12, 661-669.
- Dana, J. and Dawes, R.M. (2004). The superiority of simple alternatives to regression for social science predictions. *Journal of Educational and Behavioral Statistics*, 29(3), 317-331.
- Dangerfield, B.J. and Morris, J.S. (1992). Top-down or bottom-up: Aggregate versus disaggregate extrapolations. *International Journal of Forecasting*, 8, 233-241.
- Davies, N., Petrucci, J.D., and Pemberton, J. (1988). An automatic procedure for identification, estimation and forecasting self exciting threshold autoregressive models. *Statistician*, 37, 199-204.

- Dawes, R.M. and Corrigan, B. (1974). Linear models in decision making. *Psychological Bulletin*, 81, 95-106.
- Dawes, R.M. (1971). A case study of graduate admissions: Application of three principles of human decision making. *American Psychologist*, 26, 180-188.
- De Gooijer, J.G. and Kumar, V. (1992). Some recent developments in non-linear time series modeling, testing and forecasting. *International Journal of Forecasting*, 8, 135-156.
- De Gooijer, J.G. and Vidiella-i-Anguera, A. (2004). Forecasting threshold cointegrated systems. *International Journal of Forecasting*, 20, 237-253.
- Delbecq, A. L., Van de Ven, A. H., and Gustafson, D. H. (1975). *Group techniques for program planning*. Glenview, IL: Scott, Foresman, and Company.
- Delbecq, A.L., Van de Ven, A.H., and Gustafson, D.H. (1986). *Group techniques for program planning: A guide to Nominal Group and Delphi processes*. Middletown, WI: Greenbriar.
- Department of Navy (2001). *Gap analysis* (<http://www.doncio.navy.mil/workforce/default.htm>).
- Department of Transportation (2006). *Workforce plan 2003-2008 FY 2006 update*. Washington DC: U.S. Department of Transportation.
- Dietterich, T., Kearns, M., and Mansour, Y. (1996). Applying the weak learning framework to understand and improve C4.5. *Proceedings of the 13th International Conference on Machine Learning* (pp. 96-104). July 3-6, Bari, Italy.
- Dietterich, T.G. and Kong, E.B. (1995). Machine learning bias, statistical bias and statistical variance of decision tree algorithms. *Proceedings of the 12th International Conference on Machine Learning*. July 9-12, Tahoe City, CA.
- Dobra, A. and Gehrke, J. (2002). SECRET: A scalable linear regression tree algorithm. In *The 8th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*. July 23-26, Edmonton, Alberta, Canada.
- Dowling, K.L. and Louis, R.D. (2000). Asynchronous implementation of the nominal group technique: Is it effective? *Decision Support Systems*, 29(3), 229-248.
- Draper, B.A., Brodley, C.E., and Utgoff, P.E. (1994). Goal-directed classification using linear machine decision trees. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 16(9), 888-893.
- Dubra, E. and Gulbe, M. (2008). Forecasting the labour force demand and supply in Latvia. *Baltic Journal on Sustainability*, 14(3), 279-299.
- Duffield, J.A. and Coltrane, R. (1992). Testing for disequilibrium in the hired farm labor market. *American Journal of Agricultural Economics*, 74(2), 412-420.
- Dumpe, M.L., Herman, J., and Young, S.W. (1998). Forecasting the nursing workforce in a dynamic health care market. *Nursing Economics*, 16(4), 170-188.
- Eicher, T.S. (1996). Interaction between endogenous human capital and technological change. *Review of Economic Studies*, 63(1), 127-144.
- Elrod, T., Louviere, J.J., and Davey, K.S. (1989). *How well can compensatory models predict choice for efficient choice sets?* School of Business, Vanderbilt University, Nashville, TN.
- Esposito, F., Malerba, D., and Semeraro, G. (1997). A comparative analysis of methods for pruning decision trees. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 19(5), 476-491.
- Faché, W. (1993). The policy-developing and participative Delphi research method. In: W. Leirman (Ed.), *Delphi-Project 'Education '92': Conclusions and policy options* (pp. 106-121). Katholieke Universiteit, Leuven, Belgium.

- Faraway, J. and Chatfield, C. (1998). Time series forecasting with neural networks: A comparative study using the airline data, *Journal of the Royal Statistical Society: Series C: Applied Statistics*, 47, 231-250.
- Fayyad, U., Djorgovski, S.G., and Weir, N. (1996). Automating the analysis and cataloging of sky surveys. In U. Fayyad, G. Piatetsky-Shapiro, P. Smyth and R. Uthurusamy (Eds.), *Advances in knowledge discovery and data mining* (pp. 471-493). AAAI/MIT Press.
- Fayyad, U.M. and Irani K.B. (1992). The attribute selection problem in decision tree generation. In *Proceedings of 10th National Conference on Artificial Intelligence*. July 12-16, San Jose, CA.
- Fildes, R. (1991). Efficient use of information in the formation of subjective industry forecasts. *Journal of Forecasting*, 10, 597-617.
- Fildes, R., Goodwin, P., Lawrence, M., and Nikolopoulos, K. (2009). Effective forecasting and judgmental adjustments: An empirical evaluation and strategies for improvement in supply chain planning. *International Journal of Forecasting*, 25, 3-23.
- Fisher, D.H. and Schlimmer, J. (1988). Concept simplification and prediction accuracy. *Proceedings of the 5th International Conference on Machine Learning* (pp. 22-28). Ann Arbor, MI.
- Foster, B., Collopy, F., and Ungar, L. (1992). Neural network forecasting of short, noisy time series. *Computers and Chemical Engineering*, 16, 193-297.
- Fox, W.M. (1989). The improved Nominal Group Technique (INGT). *Journal of Management Development*, 8(1), 20-27.
- Funahashi, K. (1989). On the approximate realization of continuous mappings by neural networks. *Neural Networks*, 2(3), 183-192.
- Funke, M. (1990). Assessing the forecasting accuracy of monthly vector autoregressive models: The case of five OECD countries. *International Journal of Forecasting*, 6, 363-378.
- Gallagher, M., Hares, T., Spencer, J., Bradshaw, C., and Webb, I. (1993). The Nominal Group Technique: A research tool for general practice. *Family Practice*, 10(1), 76-81.
- Gama, J. and Brazdil, P. (1999). Linear tree. *Intelligent Data Analysis*, 3(1), 1-22.
- Ganzach, Y.A., Kluger, N., and Klayman, N. (2000). Making decisions from an interview: Expert measurement and mechanical combination. *Personnel Psychology*, 53, 1-20.
- Gates, S.M., Eibner, C., and Keating, E.G. (2006). *Civilian workforce planning in the Department of Defense: Different levels, different roles*. Santa Monica, CA: RAND Corporation.
- Georgia. Strategic planning and workforce planning, <http://www.gms.state.ga.us/agencysservices/wfplanning/index.asp>.
- Ghosh, R. (2005). A brief review of forecasting labor and skills shortages. Available at SSRN: <http://ssrn.com/abstract=870250>.
- Gooijer, J.G.D. and Hyndman, R.J. (2006). 25 years of time series forecasting. *International Journal of Forecasting*, 22, 443-473.
- Gordon, T. and Pease, A. (2006). RT Delphi: an efficient, “round-less”, almost real time Delphi method. *Technological Forecasting and Social Change*, 73(4), 321-333.
- Green, P.E. and Rao, V.R. (1971). Conjoint measurement for quantifying judgmental data. *Journal of Marketing Research*, 8(3), 355-363.
- Green, P.E. and Srinivasan, V. (1978). Conjoint analysis in consumer research: Issues and outlook. *Journal of Consumer Research*, 5(3), 103-123.

- Green, P.E. and Srinivasan, V. (1990a). *A bibliography on conjoint analysis and related methodology in marketing research*. The Wharton School, University of Pennsylvania, Philadelphia, PA.
- Green, P.E. and Srinivasan, V. (1990b). Conjoint analysis in marketing: New developments with implications for research and practice. *Journal of Marketing*, 54(4), 3-19.
- Green, P.E. (1984). Hybrid models for conjoint analysis: An expository review. *Journal of Marketing Research*, 21, 155-159.
- Green, P.E., Goldberg, S.M., and Montemayor, M. (1981). A hybrid utility estimation model for conjoint analysis. *Journal of Marketing*, 45, 33-41.
- Green, P.E., Helsen, K., and Shandler, B. (1988). Conjoint validity under alternative profile presentations. *Journal of Consumer Research*, 15, 392-397.
- Green, P.E., Krieger, A., and Wind, Y. (2001). Thirty years of conjoint analysis: reflections and prospects. *Interfaces*, 31, 56-73.
- Green, P.E., Krieger, A.M., and Agarwal, M.K. (1990). *Adaptive conjoint analysis: Some caveats and suggestions*. University of Pennsylvania, Philadelphia, PA.
- Green, T.B. (1975). An empirical analysis of nominal and interacting groups. *Academy of Management Journal*, 18, 63-73.
- Grove, W.M. and Meehal, P.E. (1996). Comparative efficiency of informal (subjective, impressionistic) and formal (mechanical, algorithmic) prediction procedures: The clinical-statistical controversy. *Psychology, Public Policy, and Law*, 2, 293-323.
- Guide To Workforce Planning at the Department of Energy (2005). *Forecasting our future demands to getting the right people; with the right skills; in the right place at the right time*. Washington DC: Human Capital Management Strategic Planning and Vision.
- Gustafson, D.H., Shukla, R.K., Delbecq, A., and Walster, G.W. (1973). A comparative study of differences in subjective likelihood estimates made by individuals, interacting groups, Delphi groups, and nominal groups. *Organizational Behavior and Human Performance*, 9, 280-291.
- Gustafsson, A., Herrmann, A., and Hube, F. (2007). Conjoint analysis as an instrument of market research practice. In A. Gustafsson, A. Herrmann, and F. Huber (Eds.), *Conjoint measurement: Methods and applications* (pp. 5-46). Berlin: Springer.
- Hagerty, M.R. (1985). Improving the predictive power of conjoint analysis: The use of factor analysis and cluster analysis. *Journal of Marketing Research*, 22, 168-184.
- Hagerty, M.R. (1986). The cost of simplifying preference models. *Marketing Science*, 5, 298-319.
- Hartmann, C.R.P., Varshney, P.K., Mehrotra, K.G., and Gerberich, C.L. (1982). Application of information theory to the construction of efficient decision trees. *IEEE Transactions on Information Theory*, IT-28, 565-577.
- Hastie, T.J. and Tibshirani, R.J. (1991). *Generalized additive models*. London: Chapman and Hall.
- Heath, D., Kasif, S., and Salzberg, S. (1993a). Induction of oblique decision trees. In *Proceedings of the 13th International Joint Conference on Artificial Intelligence* (pp. 1002-1007). August 28-September 3, Chambéry, France.
- Heath, D., Kasif, S., and Salzberg, S. (1993b). k-DT: A multi-tree learning method. In *Proceedings of the 2nd International Workshop on Multistrategy Learning* (pp.138-149). George Mason University, Harpers Ferry, WV.

- Hegarty, E.H. (1977). The problem identification phase of curriculum deliberation: Use of the Nominal Group Technique. *Curriculum Studies*, 9(1), 31-41.
- Henrich, T.R. and Greene, T.J. (1991). Using the nominal group technique to elicit roadblocks to an MRP II implementation. *Computers & Industrial Engineering*, 21(1-4), 335-338.
- Herbert, T.T. and Yost, E.B. (1979). A comparison of decision quality under nominal and interacting consensus group formats: The case of the structured problem. *Decision Sciences*, 10, 358-370.
- Hill, T., O'Connor, M., and Remus, W. (1996). Neural network models for time series forecasts. *Management Science*, 42(7), 1082-1092.
- Holliman, C., Wuerz, R., Chapman, D., and Hirshberg, A. (1997). Workforce projections for emergency medicine: how many emergency physicians does the United States need? *Academic Emergency Medicine*, 4(7), 725-730.
- Hotta, L.K. (1993). The effect of additive outliers on the estimates from aggregated and disaggregated ARIMA models. *International Journal of Forecasting*, 9, 85-93.
- Hsu, T.H. and Yang, T.H. (2000). Application of fuzzy analytic hierarchy process in the selection of advertising media. *Journal of Management and Systems*, 7 (1), pp. 19-39.
- Huber, J. and Hansen, D. (1986). Testing the impact of dimensional complexity and affective differences of paired concepts in adaptive conjoint analysis. In M. Wallendorf and P. Anderson (Eds.), *Advances in consumer research*, 14 (pp. 159-163). Provo, UT: Association for Consumer Research.
- International Personnel Management Association (IPMA-HR), (2000). *Workforce planning resource guide for public sector human resource professionals*. U.S. Department of Homeland Security, Alexandria, VA.
- International Personnel Management Association (IPMA-HR) (2002). *Workforce planning resource guide for public sector human resource professionals*. U.S. Department of Homeland Security, Alexandria, VA.
- Iowa. Demographics and Trends, <http://www.state.ia.us/idop/rfts/PERB%20Presentation%20-%20Oct%2024%202000.ppt>.
- Ishikawa, A., Amagasa, M., Shiga, T., Tomizawa, G., Tatsuta, R., and Mieno, H.R. (1993). The max-min Delphi method and fuzzy Delphi term method via fuzzy integration. *Fuzzy Sets and Systems*, 55(3), 241-253.
- Jacobs, J.M. (1996). Essential assessment criteria for physical education teacher education programs: A Delphi study (Unpublished Doctoral Dissertation). West Virginia University, Morgantown, WV.
- Janikow, C.Z. (1998). Fuzzy decision trees: Issues and methods. *IEEE Transactions on Systems, Man, and Cybernetics-Part B: Cybernetics*, 28(1), 1-14.
- Job outlook is good. (1999). *IIE Solutions*, 31(8), 14.
- John, G. (1995). Robust decision trees: Removing outliers in databases. In *Proceedings of the First International Conference on Knowledge Discovery and Data Mining* (pp. 174-179). Montreal, Canada. San Mateo, CA: AAAI Press.
- Johnson R.M. (1987). Adaptive conjoint analysis. In *Sawtooth Software Conference on Perceptual Mapping, Conjoint Analysis, and Computer Interviewing* (pp. 253-265) Ketchum, ID: Sawtooth Software.
- Johnson, R.M. (1974). Trade-off analysis of consumer values. *Journal of Marketing Research*, 11, 121-127.

- Joyce, C., McNeil, J., and Stoelwinder, J. (2006). More doctors, but not enough: Australian medical workforce supply 2001–2012. *The Medical Journal of Australia*, 184, 441-446.
- Kaastra, I. and Boyd M. (1996). Designing a neural network for forecasting financial and economic time series. *Neurocomputing*, 10, 215-236.
- Kalkanis, G. (1993). The application of confidence interval error analysis to the design of decision tree classifiers. *Pattern Recognition Letters*, 14(5), 355-361.
- Kamakura, W. (1988). A least squares procedure for benefit segmentation with conjoint experiments. *Journal of Marketing Research*, 25, 157-167.
- Kansas. Workforce Planning Report, <http://da.state.ks.us/ps/documents/00work.pdf>.
- Keel, J. (2006). *Workforce planning guide*. State Auditor's Office, Edition SAO Report No. 06-704 (<http://sao.hr.state.tx.us/Workforce/06-704.pdf>).
- Kellough, J.E. and Selden, S.C. (2003). The reinvention of public personnel administration: An analysis of the diffusion of personnel management reforms in the states. *Public Administration Review*, 63(2), 165-176.
- Kim, H. and Koehler, G.J. (1994). An investigation on the conditions of pruning an induced decision tree. *European Journal of Operational Research*, 77, 82-95.
- Kim, H. and Loh, W.Y. (2001). Classification trees with unbiased multiway splits. *Journal of the American Statistical Association*, 96, 589–604.
- Kim, J.H. (2003). Forecasting autoregressive time series with bias-corrected parameter estimators. *International Journal of Forecasting*, 19, 493–502.
- Kleinmuntz, B. (1990). Why we still use our heads instead of formulas: toward an integrative approach. *Psychological Bulletin*, 107, 296–310.
- Kohavi, R. and Li, C.H. (1995). Oblivious decision trees, graphs, and top-down pruning. In *Proceedings of the 14th International Joint Conference on Artificial Intelligence* (pp. 1071-1077). August 20-25, Montréal, Quebec, Canada.
- Kolb, R.A. and Stekler, H.O. (1992). Information content of long-term employment forecasts. *Applied Economics*, 24(6), 593-596.
- Kovner, C.T. and Reimers, C. (1998). Data sources to estimate local area supply and demand for nurses. *Public Health Nursing*, 15(2), 123-130.
- Krieger, A.M. and Green, P.E. (1988). *On the generation of Pareto Optimal, conjoint profiles from orthogonal main effects plans*. The Wharton School, University of Pennsylvania, Philadelphia, PA.
- Krolzig, H., Marcellino, M., and Mizon, G.E. (2002). A Markov-switching vector equilibrium correction model of the UK labour market. *Empirical Economics*, 27(2), 233-254.
- Kwok, S.W. and Carter, C. (1990). Multiple decision trees. In *UAI '88 Proceedings of the Fourth Annual Conference on Uncertainty in Artificial Intelligence* (pp. 327-335). Elsevier Science, Amsterdam, The Netherlands, North-Holland Publishing.
- Lago, P.P., Beruvides, M.G., Jian, J.Y., Canto, A.M., Sandoval, A., and Taraban, R. (2007). Structuring group decision making in a web-based environment by using the nominal group technique. *Computers & Industrial Engineering*, 52, 277-295.
- Lavieri, M.S. and Puterman, M.L. (2009). Optimizing nursing human resource planning in British Columbia. *Health Care Management Science*, 12(2), 119-128.
- Law, A.M. and Kelton, W.D. (1991). *Simulation modeling and analysis*. New York: McGraw-Hill Inc.
- Lawrence, M.J., O'Connor, M., and Edmundson, R.H. (2000). A field study of sales forecasting accuracy and processes. *European Journal of Operational Research*, 122(1), 151-160.

- Lederer, A.L. and Mendelow, A.L. (1986). Issues in information systems planning. *Information & Management*, 10(5), 245-254.
- Lee, S.K. (2005). On generalized multivariate decision tree by using GEE. *Computational Statistics & Data Analysis*, 49(4), 1105-1119.
- Leigh, T.W., David B.M., and Summers, J.O. (1984). Reliability and validity of conjoint analysis and self-explicated weights: A comparison. *Journal of Marketing Research*, 21, 456-462.
- LeSage, J.P. and Magura, M. (1991). Using interindustry input-output relations as a Bayesian Prior in employment forecasting models. *International Journal of Forecasting*, 7(2), 231-238.
- LeSage, J.P. (1990). Forecasting metropolitan employment using an exportbase error correction model. *Journal of Regional Science*, 30(3), 307-323.
- Levy, M., Webster, J., and Kerin, R.A. (1983). Formulating push marketing strategies: A method and application. *Journal of Marketing*, 47(Fall), 25-34.
- Li, D.T. and Dorfman, J.H. (1995). A robust approach to predicting fluctuations in state-level employment growth. *Journal of Regional Science*, 35(3), 471-484.
- Li, X.B., Sweigart, J.R., Teng, J.T.C., Donohue, J.M., Thombs, L.A., and Wang, S.M. (2003). Multivariate decision trees using linear discriminants and tabu search. *IEEE Transactions on Systems, Man, and Cybernetics-Part A*, 33(2), 194-205.
- Limère1, V. and Landeghem, H.V. (2008). Workforce planning in the banking sector: A case study. In *2008 IEEE International Conference on Industrial Engineering and Engineering Management* (pp. 670-674). December 8-10, Singapore.
- Linstone, H.A. (1978). The Delphi technique. In R.B. Fowles (Ed.), *Handbook of futures research* (pp. 273-300). Westport, CT: Greenwood.
- Linstone, H.A. (1999). *Decision making for technology executives: Using multiple perspectives to improve performance*. Norwood, MA: Artech House.
- Litterman, R.B. (1986). Forecasting with Bayesian vector autoregressions-Five years of experience. *Journal of Business and Economic Statistics*, 4, 25-38.
- Longhi, S., Nijkamp, P., Reggiani, A., and Malerhofer, E. (2005). Neural network modeling as a tool for forecasting regional employment patterns. *International Regional Science Review*, 28(3), 330-346.
- Louviere, J. (1988). *Analyzing decision making - metric conjoint analysis* (Quantitative Applications in the Social Science Series). London: Sage Publications, Inc.
- Ludwig, B. (1997). Predicting the future: Have you considered using the Delphi methodology? *Journal of Extension*, 35(5), 1-4.
- Ludwig, B.G. (1994). Internationalizing extension: An exploration of the characteristics evident in a state university extension system that achieves internationalization (Unpublished Doctoral Dissertation). Columbus, OH: The Ohio State University.
- Lurie, J., Goodman, D., and Wennberg, J. (2002). Benchmarking the future generalist workforce. *Effective Clinical Practice*, 5, 58-66.
- Lutkepohl, H. (1986). Comparison of predictors for temporally and contemporaneously aggregated time series. *International Journal of Forecasting*, 2, 461-475.
- Lutz, C. (2002). Projektion des Arbeitskräftebedarfs bis 2015: Modellrechnungen auf Basis des IAB/INFORGE-Modells (Projection of the Demand for Labour Up to 2015: Model Calculations on the Basis of the IAB/INFORGE Model), *Mitteilungen Aus Der Arbeitsmarkt-Und Berufsforschung*, 35(3), 305-326.

- MacGregor, D. (2001). Decomposition for judgmental forecasting and estimation. In J.S. Armstrong (Ed.), *Principals of forecasting: A handbook for researchers and practitioners*. Norwell, MA: Kluwer Academic Publishers.
- MacPhail, A. (2001). Nominal Group Technique: a useful method for working with young people. *British Educational Research Journal*, 27(2), 161-170.
- Mahler, J.G. (1987). Structured decision making in public organizations. *Public Administration Review*, 47(4), 336-342.
- Makridakis, S. and Hibon, M. (2000). The M3-Competition: results, conclusions and implications. *International Journal of Forecasting*, 16, 451-476.
- Makridakis, S., Wheelwright, S. C., and Hyndman, R. J. (1998). *Forecasting methods for management* (3rd ed.). New York: John Wiley.
- Mangasarian, O.L. (1997). Mathematical programming in data mining. *Data Mining & Knowledge Discovery*, 1(2), 183-201.
- Marquez, L.O. (1992). *Function approximation using neural networks: A simulation study* (Ph.D. Dissertation). Honolulu, HI: University of Hawaii.
- Martorelli, W.P. (1981). Cowboy DP scouting avoids personal fumbles. *Information Systems News*, (November 16).
- Matarazzo, J.M. (2000). Library human resources: The Y2K plus 10 challenge. *Journal of Academic Librarianship*, 26(4), 223-224.
- Maynard, A. (2006). Medical workforce planning: Some forecasting challenges. *The Australian Economic Review*, 39(3), 323-329.
- Meade, N. (2000). A note on the robust trend and ARARMA methodologies used in the M3 competition. *International Journal of Forecasting*, 16, 517-519.
- Meehan, R.H. and Ahmed, S.B. (1990). Forecasting human resources requirements: A demand model. *Human Resource Planning*, 13(4), 297-307.
- Miller, L.E. (2006). Determining what could/should be: The Delphi technique and its application. Paper presented at the *2006 Annual Meeting of the Mid-Western Educational Research Association*. October, 2006, Columbus, OH.
- Moore, W.L. and Holbrook, M.B. (1990). Conjoint analysis on objects with environmentally correlated attributes: The questionable importance of representative design. *Journal of Consumer Research*, 16, 490-497.
- Moore, W.L. and Semenik, R.J. (1988). Measuring preference with hybrid conjoint analysis: The impact of a different number of attributes in the master design *Journal of Business Research*, 16, 261-274.
- Moore, C.M. (1987). *Group techniques for idea building*. Newbury Park, CA: Sage Publishing, Inc.
- Moore, C.M., (1990). *Group techniques for idea building* (2nd ed.), (Applied Social Research Methods Series). Newbury Park, CA: Sage Publications, Inc.
- Morris, W.T. (1979). *Implementation strategies for industrial engineers*. Columbus, OH: Grid Publishing Company.
- Murphy, P.M. and Pazzani, M. (1994). Exploring the decision forest: An empirical investigation of Occam's razor in decision tree induction. *Journal of Artificial Intelligence Research*, 1, 257-275.
- Murry, T.J., Pipino, L.L., and Gigch, J.P. (1985). A pilot study of fuzzy set modification of Delphi. *Human Systems Management*, 5(1), 76-80.

- Murthy, S.K., Kasif, S., and Salzberg, S. (1994). A system for induction of oblique decision trees. *Journal of Artificial Intelligence Research*, 2, 1-32.
- National Wildlife Refuge System (2008). *Strategic workforce planning report*. Washington, DC.
- Nelson, J.S., Jayanthi, M., Brittain, C.S., Epstein, M.H., and Bursuck, W.D. (2002). Using the Nominal Group Technique for homework communication decisions: An exploratory study. *Remedial and Special Education*, 23(6), 380-387.
- Nelson, M., Hill, T., Remus, W., and O'Connor, M. (1999). Time series forecasting using neural networks: Should the data be deseasonalized first? *Journal of Forecasting*, 18, 359-370.
- Nemiroff, P.M., Pasmore, W.A., and Ford, D. L. (1976). The effects of two normative structural interventions on established and ad hoc groups: Implications for improving decision making effectiveness. *Decision Sciences*, 7, 841-855.
- Newton, J.R. (1965). Judgment and feedback in a quasi-clinical situation. *Journal of Personality and Social Psychology*, 1, 336-342.
- O'Connor, M., Remus, W., and Lim, K. (2005). Improving judgmental forecasts with judgmental bootstrapping and task feedback support. *Journal of Behavioral Decision Making*, 18, 247-260.
- O'Brien-Pallas, L., Baumann, A., Donner, G., Murphy, G.T., Lochhaas-Gerlach, J., and Luba, M. (2001). Forecasting models for human resources in health care. *Journal of Advanced Nursing*, 33, 120-129.
- Oliveira A.L. and Sangiovanni-Vincentelli, A. (1995). Using the minimum description length principle to infer reduced order decision graphs. *Machine Learning*, 26, 23-50.
- Oliver, J.J. (1993). Decision graphs - an extension of decision trees. In *Proceedings of the 4th International Conference on Artificial Intelligence and Statistics* (pp. 343-350). Fort Lauderdale, FL.
- Park, H.A. and Choi, E. (2001). Projected supply and demand for nurses in Korea by 2015. *Journal of Nursing Scholarship*, 33(4), 387.
- Parzen, E. (1982). ARARMA models for time series analysis and forecasting. *Journal of Forecasting*, 1, 67-82.
- Pazzani, M., Merz, C., Murphy, P., Ali, K., Hume, T., and Brunk, C. (1994). Reducing misclassification costs. In *Proceedings of the 11th International Conference of Machine Learning* (pp. 217-225). July 10-13, New Brunswick, NJ.
- Pedrycz, W. and Sosnowski, Z.A. (2000). Designing decision trees with the use of fuzzy granulation. *IEEE Transactions on Systems, Man, and Cybernetics-Part A*, 30(2), 151-159.
- Pedrycz, W. and Sosnowski, Z.A. (2005). C-fuzzy decision trees. *IEEE Transactions on Systems, Man, and Cybernetics-Part C*, 35(4), 498-511.
- Pemberton, J. (1989). *Forecast accuracy of nonlinear time series models* (Technical Report 77). Statistics Research Center, Graduate School of Business, University of Chicago, Chicago, IL.
- Peng Y.H., Chen T.J., and Tse, P.W. (2000). Fuzzy inductive learning in data mining and its application to machine fault diagnosis. *2000 International Conference on Advanced Manufacturing Systems and Manufacturing Automation* (pp. 158-163). Guangzhou, China.
- Persad, K.R., O'Connor, J.T., and Varghese, K. (1995). Forecasting engineering manpower requirements for highway preconstruction activities. *Journal of Management in Engineering*, 11(3), 41-47.

- Persaud, D., Cockerill, R., Pink, G., and Trope, G. (1999). Determining Ontario's supply and requirements for ophthalmologists in 2000 and 2005: A comparison of projected supply and requirements. *Canadian Journal of Ophthalmology*, 34, 82-87.
- Peter, W. and Kriewall, M.A. (1980). State of mind effects on the accuracy with which utility functions predict marketplace choice. *Journal of Marketing Research*, 19, 277-293.
- Phillips, R. (2000). New applications for Delphi technique. *Annual V.2: Consulting*, 2, 191-195.
- Popkov, A.Y., Popkov, Y.S., and Van Wissen, L. (2005). Positive dynamic systems with entropy operator: Application to labour market modeling. *European Journal of Operational Research*, 164(3), 811-828.
- Prescott, P.A. (1991). Forecasting requirements for health care personnel. *Nursing Economics*, 9(1), 18-24.
- Quenouille, M.H. (1957). *The analysis of multiple time-series* (2nd ed., 1968). London: Griffin.
- Quinlan, J.R. (1986). Induction of decision trees. *Machine Learning*, 1, 81-106.
- Quinlan, J.R. (1987). Simplifying decision trees. *International Journal of Man-Machine Studies*, 27, 221-234.
- Quinlan, J.R. (1993). *C4.5: Programs for machine learning*. San Francisco, CA: Morgan Kaufmann.
- Rauch, W. (1979). The decision Delphi. *Technological Forecasting & Social Change*, 15, 159-169.
- Reilly, R.R. and Chao, G.T. (1982). Validity and fairness of some alternative employee selection procedures. *Personnel Psychology*, 35, 1-62.
- Reisman, S., Johnson, T.W., and Mayes, B.T. (1992). Group decision program: A videodisc-based group decision support system. *Decision Support Systems*, 8(2), 169-180.
- Roberfroid, D., Leonard, C., and Stordeur, S. (2009). Physician supply forecast: better than peering in a crystal ball? *Human Resources for Health*, 7(10), 1-13.
- Roberfroid, D., Stordeur, S., Camberlin, C., Van de Voorde, C., Vrijens, F., and Léonard, C. (2008). *Physician workforce supply in Belgium: current situation and challenges*. Brussels, Belgium: Belgian Health Care Knowledge Centre.
- Roig, A.H. (1999). Testing Spanish labour market segmentation: An unknown regime approach. *Applied Economics*, 31(3), 293-305.
- Rosenfeld, Y. and Warszawski, A. (1993). Forecasting methodology of national demand for construction labour. *Construction Management & Economics*, 11(1), 18-29.
- Rowe, G. and Wright, G. (1999). The Delphi technique as a forecasting tool: Issues and analysis. *International Journal of Forecasting*, 15, 353-375.
- Rowe, G., Wright G., and Bolger, F. (1991). Delphi: A reevaluation of research and theory. *Technological Forecasting & Social Change*, 39, 235-251.
- Rumelhart, D.E., McClelland, J.L., and the PDP Research Group (1986). *Parallel distributed processing: Explorations in the microstructure of cognition (Volume I)*. Cambridge, MA: MIT Press.
- Sanders, N.R. and Ritzman, L.P. (2001). Judgmental adjustment of statistical forecasts. In J.S. Armstrong (Ed.), *Principles of forecasting: A handbook for researchers and practitioners* (pp. 405-416). Norwell, MA: Kluwer Academic Publishers.
- Sarantis, N. (2001). Nonlinearities, cyclical behaviour and predictability in stock markets: International evidence. *International Journal of Forecasting*, 17, 459-482.
- Saremi, M., Mousavi, S.F., and Sanayei, A. (2009). TQM consultant selection in SMEs with TOPSIS under fuzzy environment. *Expert Systems with Applications*, 36(2), 2742-2749.

- Schwartz, D. (1978). Locked box combines survey methods, helps end woes of probing industrial field. *Marketing News* (January 27), 18.
- Selden, S.C., Ingraham, P., and Jacobson, W. (2001). Human resource practices in state government: Findings from a national survey. *Public Administration Review*, 61(5), 598-607.
- Setiono, R. and Liu, H. (1999). A connectionist approach to generating oblique decision trees. *IEEE Transactions on Systems, Man, and Cybernetics-Part B*, 29, 440-444.
- Sexton, R.S., Dorsey, R.E., and Johnson, J.D. (1998). Towards global optimization of neural networks: A comparison of the genetic algorithms and backpropagation. *Decision Support Systems*, 22, 171-185.
- Sharda, R. and Patil, R.B. (1992). Connectionist approach to time series prediction: An empirical test. *Journal of Intelligent Manufacturing*, 3, 317-323.
- Shipman, S., Lurie, J., and Goodman, D. (2004). The general pediatrician: Projecting future workforce supply and requirements. *Pediatrics*, 113, 435-442.
- Shlien, S. (1992). Nonparametric classification using matched binary decision trees. *Pattern Recognition Letters*, 13(2), 83-88.
- Shukla, S.K. and Tiwari, M.K. (2009). Soft decision trees: A genetically optimized cluster oriented approach. *Expert Systems with Applications*, 36(1), 551-563.
- Shukla, S.K. (2009). Decision models and artificial intelligence in supporting workforce forecasting and planning (Master's Thesis). University of Texas at San Antonio, TX.
- Smith, S.K. and Sincich, T. (1990). On the relationship between length of base period and population forecast errors. *Journal of the American Statistical Association*, 85, 367-375.
- Smyth, P., Gray, A., and Fayyad, U.M. (1995). Retrofitting decision tree classifiers using kernel density estimation. In *Proceedings of the 12th International Conference on Machine Learning* (pp. 506-514). July 9-12, Tahoe City, CA.
- Sorensen, K. and Janssens, G.K. (2003). Data mining with genetic algorithm on binary trees. *European Journal of Operational Research*, 151, 253-264.
- Srinivasan, V. and Shocker, A.D. (1973). Estimating the weights for multiple attributes in a composite criterion using pair-wise judgments. *Psychometrika*, 38, 473-493.
- Srinivasan, V., Jain, A.K., and Malhotra, N.K. (1983). Improving predictive power of conjoint analysis by constrained parameter estimation. *Journal of Marketing Research*, 20, 433-438.
- State of Washington (2000). *Workforce planning guide: Right people, right jobs, right time*. Olympia, WA.
- Steckel, J.H., DeSarbo, S.W., and Mahajan, V. (1990). On the creation of feasible conjoint analysis experimental designs. *Decision Sciences*, 21, 295-315.
- Stewart, T.R. (2001). Improving reliability of judgmental forecasts. In J.S. Armstrong (Ed.), *Principles of forecasting: A handbook for researchers and practitioners* (pp. 81-106). Norwood, MA: Kluwer Academic Publishers.
- Sweeney, S.H. (2004). Regional occupational employment projections: Modeling supply constraints in the direct-requirements approach. *Journal of Regional Science*, 44(2), 263-288.
- Tape, T.G., Kripal, J., and Wigton, R.S. (1992). Comparing methods of learning clinical prediction from case simulation. *Medical Decision Making*, 12(3), 213-221.
- Tarlov, A.R. (1995). Estimating physician workforce requirements: The devil is in the assumptions. *JAMA: Journal of the American Medical Association*, 274(19), 1558-1560.

- Tashman, L. and Hoover, J. (2001). Diffusion of forecasting principles: An assessment of forecasting software programs. In J.S. Armstrong (Ed.), *Principles of forecasting: A handbook for researchers and practitioners*. Norwell, MA: Kluwer Academic Publishers.
- Thompson, J.R. and Mastracci, S.H. (2005). *The blended workforce: Maximizing agility through nonstandard work arrangements*. Washington, DC: IBM Center for the Business of Government.
- Tong, H. (1983). *Threshold models in non-linear time series analysis*. New York: Springer-Verlag.
- Tong, H. (1990). *Non-linear time series: A dynamical system approach*. Oxford: Clarendon Press.
- Tseng, K.H., Lou, S.J., Diez, C.R., and Yang, H.J. (2006). Using online Nominal Group Technique to implement knowledge transfer. *Journal of Engineering Education*, 95(4), 335.
- Tuffrey-Wijne I., Bernal J., Butler G., Hollins S., and Curfs, L. (2007). Using Nominal Group Technique to investigate the views of people with intellectual disabilities on end of-life care provision. *Journal of Advanced Nursing*, 58(1), 80-89.
- Turoff, M. and Hiltz, S.R. (1996). Computer based Delphi process. In M. Adler, & E. Ziglio (Eds.), *Gazing into the oracle: The Delphi method and its application to social policy and public health* (pp. 56-88). London, UK: Jessica Kingsley Publishers.
- Ulschak, F.L. (1983). *Human resource development: The theory and practice of need assessment*. Reston, VA: Reston Publishing Company, Inc.
- Utgoff, P.E. and Brodley, C.E. (1990). An incremental method for finding multivariate splits for decision trees. In *Proceedings of the 7th International Conference on Machine Learning* (pp. 58-65). June 21-23, Austin, TX.
- Utgoff, P.E. and Clouse, J.A. (1996). *A Kolmogorov-Smirnoff metric for decision tree induction* (Technical Report 93-3). Department of Computer Science, University of Massachusetts, Amherst, MA.
- Utgoff, P.E. (1989). Incremental induction of decision trees. *Machine Learning*, 4, 161-186.
- Van der Laan, L. (1996). A review of regional labour supply and demand forecasting in the European Union. *Environment & Planning A*, 28(12), 2105-2123.
- Volterra, V. (1930). *Theory of functionals and of integral and integro-differential equations*. New York: Dover.
- Wang, X., Chen, B., Qian G., and Ye, F. (2000). On the optimization of fuzzy decision trees. *Fuzzy Sets and Systems*, 112, 117-125.
- Ward, D. (1996). Workforce demand forecasting techniques. *Human Resource Planning*, 19(1), 54-55.
- Weber, R. (1992). Fuzzy ID3: A class of methods for automatic knowledge acquisition. In *Proceedings of the 2nd International Conference on Fuzzy Logic Neural Networks* (pp. 265-268). Iizuka, Japan.
- White, A.P. and Liu, W.Z. (1994). Bias in information-based measures in decision tree induction. *Machine Learning*, 15, 321-329.
- Wiener, N. (1958). *Non-linear problems in random theory*. London: Wiley.
- Wild, C. and Torgersen, H. (2000). Foresight in medicine: lessons from three European Delphi studies. *European Journal of Public Health*, 10(2), 114-119.

- Wiley, J.B. (1977). Selecting Pareto Optimal subsets from multiattribute alternatives. In W.D. Perreault, Jr. (Ed.), *Advances in Consumer Research* (Vol. 4, pp. 171-174). Atlanta, GA: Association for Consumer Research.
- Willems, E.J.T.A. and de Grip, A., (1993). Forecasting replacement demand by occupation and education. *International Journal of Forecasting*, 9(2), 173-185.
- Winston, W.L. (1994). *Operations research: Applications and algorithms*. Pacific Grove, CA: Duxbury Press.
- Wirth, J. and Catlett, J. (1988). Experiments on the costs and benefits of windowing in ID3. In *Proceedings of the 5th International Conference on Machine Learning* (pp. 87-99). June 12-14, University of Michigan, Ann Arbor, MI.
- Wittink, D. and Cattin, P. (1989). Commercial use of conjoint analysis: An update. *Journal of Marketing*, 53(3), 91-96.
- Wittink, D., Vriens, M., and Burhenne, W. (1994). Commercial use of conjoint analysis in Europe: Results and critical reflections. *International Journal of Research in Marketing*, 11, 41-52.
- Wittink, D.R. and Bergestuen, T. (2001). Forecasting with conjoint analysis. In J.S. Armstrong (Ed.), *Principles of forecasting: A handbook for researchers and practitioners*. Norwell, MA: Kluwer Academic Publishers.
- Wright, P. and Kriewall, M.A. (1980). State of mind effects on the accuracy with which utility functions predict marketplace choice. *Journal of Marketing Research*, 17, 277-293.
- Yildiz, O.T. and Alpaydin, E. (2001). Omnivariate decision trees. *IEEE Transactions on Neural Networks*, 12(6), 1539-1546.
- Zellner, A. (1971). *An introduction to Bayesian inference in econometrics*. New York: Wiley.
- Zhang, G.B., Patuwo B.E., and Hu, M.Y. (1998). Forecasting with artificial neural networks: The state of the art. *International Journal of Forecasting*, 14, 35-62.
- Zheng, Z. (1995). Constructing nominal X-of-N attributes. In *Proceedings of the 14th International Joint Conference on Artificial Intelligence* (pp. 1064-1070). August 20-25, Montreal, Quebec, Canada.
- Zipfinger, S. (2007). *Computer aided Delphi: An experimental study of comparing round-based with real-time implementation of the method*. Linz, Austria: Trauner Verlag.
- Zolingen, S.J. and Klaassen, C.A. (2003). Selection processes in a Delphi study about key qualifications in senior secondary vocational education. *Technological Forecasting & Social Changes*, 70, 317-340.
- Zurn, P., Dal Poz, M., Stilwell, B., and Adams, O. (2004). Imbalance in the health workforce. *Human Resources for Health*, 2(1), 1-12.

PART II – DECISION SUPPORT SOFTWARE PROGRAM USING ARTIFICIAL INTELLIGENCE TECHNIQUES FOR WORKFORCE FORECASTING

8.0 SUMMARY

This project studies workforce forecasting in two main aspects: (1) an extensive review of the existing methodologies and techniques and (2) an effort to develop a decision support system with models and software programming. The main objective of this project is to enhance the effectiveness of workforce forecasting and deployment through innovative approaches of artificial intelligence. The deliverables include a summary report of conventional and innovative methods of workforce forecasting (Part I) and a decision support software program using artificial intelligence techniques for workforce forecasting (Part II).

Part II covers the development of a forecasting decision support system. This system consists of a forecasting model powered by artificial intelligence, a data collection scheme that covers a full spectrum of conditions, a software program to realize the proposed model, and simulated cases. In this part of the project, the research team attempts to achieve the following targets: (1) develops an iterative and incremental process for the workforce forecasting, (2) integrates the strategic business goals with workforce planning process, (3) suggests a longer planning horizon for accurate forecasts, (4) implements a broad domain of workforce planning activities, from environmental scanning to strategic reviews, (5) reduces the future uncertainty by incorporating the skill's gap analysis, and (6) provides a competency framework which supports coherent analysis of several factors involved.

As an integrated solution, a hybrid model has been proposed, which integrates various artificial intelligent soft computing techniques, including ant-based algorithm, genetic algorithm, and neural network, and fuzzy logic. A software program has been developed to realize the proposed hybrid model with a user-friendly graphical interface. Some test results show that the hybrid model performs well while solving simulated cases.

Following is a summary of the main contents and contributions of Part II:

- This report presents a more comprehensive set of questionnaire with 52 questions and 17 parameters to capture the human resource information within most organizations.
- This report presents an improved intelligent decision support system for workforce forecasting. After conducting several tests on various artificial intelligence techniques, a “*Self Guided Ant-based Genetically-Optimized-Neural-Network*” (SGA GONN) model is proposed, which is a robust solution methodology that outperforms the existing solution methodologies compared in this research.
- A software program has been developed that integrates MATLAB, MS Excel, and Visual C#. It has a graphical interface that allows customization of the algorithm and also helps in maintaining the database. Users with little knowledge the models will still be able to operate the software by simply following the instructions with suggested default settings.

Part I presents a thorough literature review of fundamental research and practices of demand and supply forecasting techniques for workforce analysis based on 289 literature resources reviewed by the research team.

9.0 INTRODUCTION TO WORKFORCE FORECASTING

9.1 Background

Workforce planning is one of the crucial human resource (HR) functions for an organization to ensure that its long term strategic goals are met. Workforce planning is defined as organized process of identifying the number and mix of employees along with the types of skill sets required by them to accomplish the organization's strategic goals and objectives. It starts with a well directed strategic plan, reliable and structured workforce data, a strong internal and external environmental scanning, and a keen awareness of trends (Cotten, 2007). A successful workforce analysis plan is based on the lean principle of continuous improvement. It is an iterative process which includes four steps in succession and thrives on the underlying idea that there is always a scope for more improvement. These four steps are (Keel, 2006):

- 1) Defining organization's strategic goals and objectives,
- 2) Conducting workforce analysis,
- 3) Implementation of workforce plan, and
- 4) Monitoring, evaluating and revising strategies.

Workforce analysis is a most vital and complex phase in workforce planning. This phase comprises demand analysis, supply analysis, gap analysis, and strategy development. The thrust of this thesis is the demand analysis phase (i.e., forecasting of future workforce) of workforce analysis. Forecasting the future workforce demand (first phase of workforce analysis) of an organization begins with the mission area analysis of its current products, processes, policies and anticipating needed changes in the workforce. The Demand Forecast describes what each mission workforce should be (e.g., 1 to 5 years) in terms of: Size – total number of positions needed, Composition – mix of full time and part time positions, and Job requirements – competencies required by crucial work aligned to core processes.

9.2 Problem Statement and Objectives

Regardless of the methods used, it is almost impossible to forecast the exact amount of the future workforce in long planning horizons. There will always be an element of uncertainty until the forecast horizon has come to pass. Researchers and practitioners have put ample efforts in developing forecasting methods to minimize the element of uncertainty in forecasting. An intensive scanning of the literature shows that there is not much work regarding review of workforce forecasting methods. It is also seen that organizations feel the difficulty in selecting workforce forecasting methods according to their strategic direction and availability of data. There can be one or a combination of more than two suitable forecasting methods in a particular situation. Therefore, one of the biggest questions is how to choose the best forecasting method or combination of methods for forecasting future workforce? For any organization, just to know the best suitable forecasting technique(s) is not enough for effective workforce forecasting. Identification of key factors affecting a future workforce, questionnaires development for effective data collection, interpretation of responses to questionnaires, integration of responses to get final data are other vital tasks that should be performed in an organized way for effective and efficient workforce forecasting. Moreover, in the literature, decision trees are not used much for

the forecasting of future workforce; while, it is well proved that decision trees are one the best decision making techniques in the area of machine learning. This thesis answers all the questions starting from doing literature review of workforce forecasting techniques to development of a new decision tree for workforce forecasting.

9.3 Detailed Overview

There are four key steps to the workforce analysis phase of the planning model. These steps are pictorially shown in Figure 8 (Keel, 2006) followed by details of these steps.

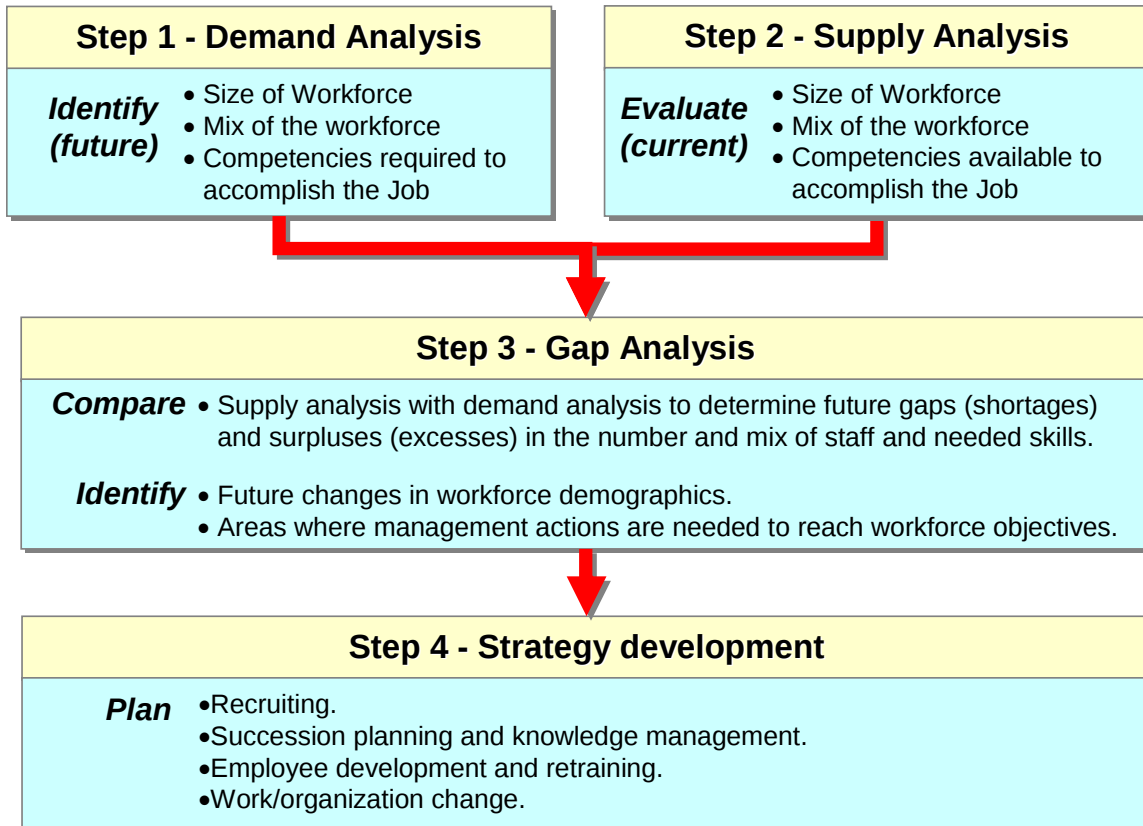


Figure 8: Four Key Steps for Conducting Workforce Analysis

9.3.1 Step 1: Demand Analysis

Forecasting the future workforce demand begins with the mission area analysis of the current and future products, processes, and organization's policies. The Demand Forecast will describe what should be the workforce size, compositions, and competencies in the planning horizon. In light of strategic direction and through collaboration with mission area experts, the true workforce requirement needed to accomplish future mission is assessed in demand analysis.

Acknowledging the difficulty of describing the unknown, a disciplined balance of potential technological, threat, product, political, and economic changes with risk assessment is needed. The demand analysis process should examine not only what work the organization will do in the future and who will do it, but how that work will be performed? Some possible considerations include:

- How will jobs and workload change as a result of technological advancements, economic, social, and political conditions?
- What are the consequences or results of these changes?
- What will be the reporting relationships?
- How will divisions, work groups, and jobs be designed?
- How will work flow into each part of the organization? What will be done with it? Where will the work flow?

Once the ‘what and how’ of future work has been determined through the identification of actual positions that are crucial to the core processes of the organization, the next step is to associate critical competencies that will be needed to carry out that work. The future workforce profile created through the demand forecast analysis will display a set of competencies that describes the future workforce.

9.3.2 Step 2: Supply Analysis

Supply analysis assesses the available workforce size, composition, and competencies according to current capabilities and policies in the planning horizon. In order to conduct this analysis, organizations should consider workforce, workload, and competencies as integrated elements. The demographic data of the organization provides snapshots of the current workforce for the supply analysis process. Identifying employment trends is one of the major factors in projecting the future workforce supply. Organizations generally use transaction data to identify employment trends. Required transaction data can be collected by reviewing changes in workforce demographics by: Occupation, Grade level, Organizational structure, Race/national origin, Gender, Age, Length of service, and Retirement eligibility.

This data also helps in developing valuable information on areas such as retirement eligibility or turnover for a given point in the future by projecting from current workforce demographic data. Personnel transaction data also provides the basis to identify baselines such as turnover rates. Moreover, it can also provide powerful tools to forecast workforce changes in the future that may occur from actions such as resignations and retirements. In conjunction with demographic data, transaction data help human resource professionals and other managers to forecast opportunities for workforce change that can be incorporated into the action plans.

When modeling the current workforce, organizations must include permanent employees, supplemental direct-hire employees, and contract workers (Cotton, 2007). Permanent employees are on the organization’s payroll, have regular work hours, and are entitled to receive the benefits of regular employment offered by the organization (Thompson and Mastracci, 2005). Supplemental direct-hire employees are on the payroll of the organization, but do not have regular work hours and are not entitled to full benefits of employment. Supplemental employees work on a temporary, seasonal, or on-call basis (Thompson and Mastracci, 2005). For examining the current supply organizations conduct SWOT (Strength, Weakness, Opportunity, and Threat) analysis. A SWOT analysis brings together the external and internal information to develop strategies and objectives. The SWOT analysis develops strategies that align organization strengths with external opportunities, identifies internal weaknesses, and acknowledges threats that could affect organization success. The SWOT analysis can be conducted by doing internal and external environmental scanning (IPMA–HR 2000). Environmental scanning leads to an

adaptive workforce plan in response to rapid workplace changes. Such scanning enables the managers to review and analyze internal and external Strengths, Weaknesses, Opportunities and Threats. Environmental scanning addresses external and internal factors that will affect short-term and long-term goals.

Workforce supply is, at its most basic level, the current workforce plus new hires less projected separations at some specific date in the future (Figure 9). For some organizations, projected workforce supply will be the result of a sophisticated mathematical model. For others, it will be an educated guess based upon data collected in the environmental scan. For most, it will be somewhere in between. No matter the level of sophistication, all models need to consider the same elements when projecting the future workforce supply. These elements are the inventory of the current workforce, the rate at which employees in specific occupations and at various leadership levels will leave the organization, and types of skills and abilities the organization will be able to attract.

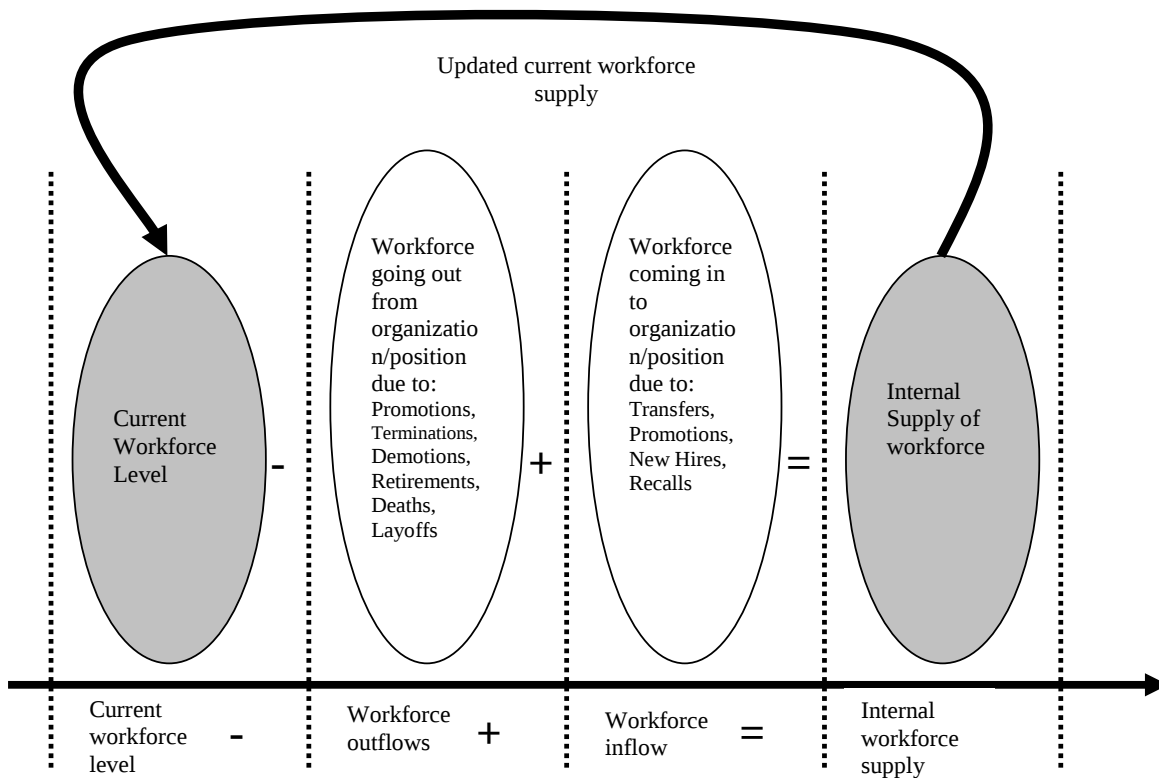


Figure 9: Workforce Supply Analysis Model

9.3.3 Step 3: Gap Analysis

Gap analysis involves comparing the results of supply and demand analysis to identify gaps between the supply and demand of total workforce and their mix and competencies. In gap analysis, managers use workload and workforce data and the competency sets developed earlier in two phases (i.e., the supply and demand analysis phases). The most important thing in conducting gap analysis is the coordination in supply and demand data and competencies analyses as they have to be comparable.

Gap analysis identifies situations when demand exceeds supply, such as when critical work demand, number of personnel, or current/future competencies will not meet future needs. It also identifies situations when future supply exceeds demand. In either event, organizations must identify these differences and make plans to address them. Depending upon how the supply and demand are determined and how specific they are gaps can be identified by job title, series, grades, and locations. To be effective for comparison, the data and competencies in the supply and demand analysis phases need to be developed in tandem. The solution analysis that will close the gaps must be strategic in nature. When doing solution analysis, organizations should be prepared to address ongoing as well as unplanned changes in the workforce. The trends identified in supply and demand analysis can help organizations to anticipate these changes.

In summary, calculating gaps will enable the organizations to identify where human capital (people and competencies) will not meet future needs (demand will exceed supply). The gap analysis also will determine whether human capital exceeds the needs of the future (supply will exceed demand). There may also be situations where supply will meet future demand, thus resulting in a zero difference or no gap. The gap analysis process is outlined in Table 3 (IPMA–HR 2000). Once the gaps between the demand and the supply are identified, it is necessary to consult with management to set priorities to fill the gaps which will have a critical impact on organization’s strategic goals.

Table 3: The Gap Analysis Process

How	What
Assess	The current supply of human capital
Factor in	Variables and assumptions
To come up with	Supply of human capital, then
Compare to	Demand
To come up with	Gaps and surpluses

9.3.4 Step 4: Strategy Development

The final step in the workforce analysis phase is development of strategies to address future gaps and surpluses. There is a wide range and combination of strategies that might be used to attract and develop the workforce with needed competencies, or to deal with excesses in competencies no longer needed for mission accomplishment. There is also numerous factors that will influence the selection of strategy or, more likely, which combination of strategies should be used. Some of these factors include, but are not limited to, the following:

- **Time-** Is there enough time to develop the workforce internally for anticipated vacancies or new competency needs, or fast-paced recruitment is the best approach?
- **Resources-** The availability of adequate resources will likely influence which strategies are used and to what degree, as well as priorities and timing.
- **Workplace and workforce dynamics-** Whether particular productivity and retention strategies need to be deployed will be influenced by workplace climate (e.g., employee satisfaction levels), workforce age, diversity, personal needs, etc.
- **Job classifications-** Do the presently used job classifications and position descriptions reflect the future functional requirements and competencies needed? Does the structure of

the classification series have enough flexibility to recognize competency growth and employee succession in a timely fashion? Does it allow compensation flexibility?

Multiple options or solution sets for filling any gaps should be developed and priced. In this way, which gaps to resolve and which solutions to fund can be selected strategically based on mission priorities, expected ROI (Return on Investment), and direct alignment to strategy.

In order to accomplish above mentioned four key steps for workforce analysis, various methodologies and approaches had been developed and can be found in literature. Since Demand Analysis is the major thrust of this thesis, in the subsequent section of this thesis we will only focus on the review of existing methodologies for conducting the demand analysis, i.e., workforce demand forecasting.

9.4 Solutions Outline

Due to the complexity and the fuzzy nature of workforce planning in the real world, a crisp mathematical formulation for workforce forecasting fails to capture a generic scenario that can be implemented for all the organizations in general (Tripathi et al., 2010). In this respect, this research proposes a new forecasting decision support system (FDSS) which can be utilized to forecast the workforce for generic scenarios of most organizations. The suggested methodology iteratively adapts itself (over the years) in accordance with the organization under consideration. The proposed methodology has its roots in neural network, adaptive genetic algorithm and algorithm of self guided ants. In general, artificial neural networks have been proved to be universal approximators and a three-layer feedforward neural network can approximate any nonlinear continuous function to a considerable accuracy (Brown and Harris, 1994). Owing to its structure, a neural network utilizes several learning algorithms for enhancing its prediction capabilities, e.g., Genetic Algorithm (GA) (Tripathi et al., 2008) and backpropagation (Pham and Karaboga, 2000). This research utilizes GA as learning algorithm and enhances the learning capability of the network by incorporating the concept of self guided ants to enhance the learning procedure. Thereafter a simulated case study of a hypothetical organization is constructed and the efficiency of the proposed algorithm is tested against a previously developed methodology Clonal C-Fuzzy Decision Tree (C²FDT) (Shukla et al., 2009) for prediction purposes.

10.0 THE FORECASTING DECISION SUPPORT SYSTEM (FDSS) FORMULATION

The problem of workforce forecasting has become a burning platform for most organizations. According to a survey on 200 organizations globally, only 14% of them are prepared to a very large extent and 58% are prepared to some extent to meet with the challenges put forth by the workforce planning activities (Tripathi et al., 2010). The loss of skills, corporate knowledge and human resource management data is the governing reason for poor operational resource planning/workforce forecasting within an organization. The incompetency of the organizations to conduct formal, integrated workforce planning plays a key role for the loss of this information. At the same time, the firms dedicated to more robust workforce planning activities have a longer time horizons for their workforce plans. Such organizations keep themselves prepared for any potential loss of the aforementioned knowledge by maintaining a structured database. Such a structured database can be effectively utilized for successfully carrying on the workforce planning activities with fewer challenges and more accurate results.

The major modeling challenges for the organizations involved in workforce planning are shown in Figure 10. The figure shows the concept of ‘5 Why’ from the *Lean Thinking* (Womack and Jones, 2003) to arrive at the final answer.

- Can we formulate a mathematical model to solve the workforce forecasting problem?
- Can a generic mathematical model be formulated to represent all the scenarios for all the organizations?
- If the generic mathematical model cannot be formulated, then can a non mathematical model be formulated for this problem type?
- How can such a generic model be implemented to meet with specific organizations context?

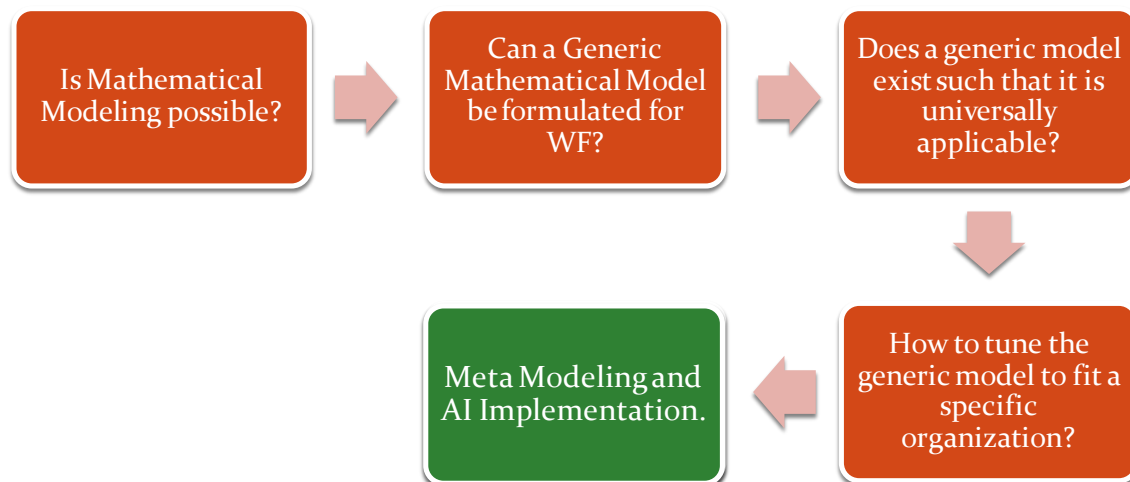


Figure 10: Modeling Challenges in Workforce Forecasting

Therefore finally we arrive to the solution that by utilizing meta-modeling and artificial intelligence implementation a generic model can be formulated and can be customized for a specific organization.

Firms support a range of approaches to workforce planning depending upon their consistency with the maturity and size of individual business units. In this respect, the proposed FDSS identifies the intrinsic relationships existing among the factors that critically affect the workforce planning and forecasting. The four essential phases which form the building blocks of the FDSS are shown in Figure 11:

- **Data Collection:** This phase involves collecting organizational data in the form of questionnaire/survey forms.
- **Parameter Identification:** Categorizing them into a generalized set of parameters.
- **Parameter Integration:** A fuzzy logic based internal engine is utilized to convert the verbal answers from the questionnaire into the software understandable set of quantified parameters.
- **A Data Mining Algorithm:** This is utilized for analyzing the hidden pattern within the data collected for predicting the future accordingly.

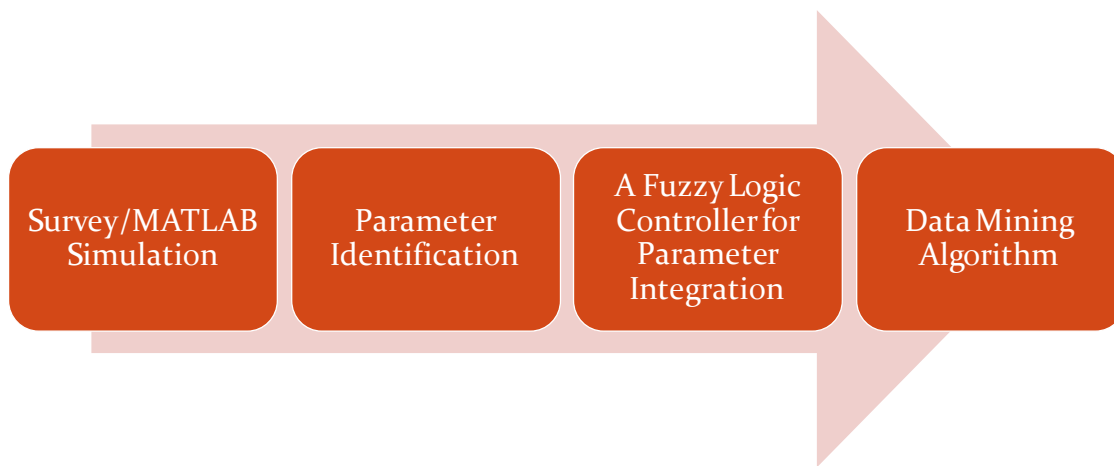


Figure 11: Building Blocks of FDSS

10.1 Data Collection

To get a deeper insight of human resource data within an organization, conducting a formal and integrated workforce planning survey is critical. However, the organization must show its active participation to the workforce analyst. Such participation includes revealing information regarding a comprehensive set of metrics for various workforce planning activities. In this respect, development of a survey forms as a follow up to the strategic workforce planning constitute an essential part of the FDSS. This research proposes an improved survey form that consists of a comprehensive set of 52 questions and has been designed to gain the complete insight of an organization's HR information (a comprehensive list of all the questions is presented in Table 4). Moreover, for each question in the survey form, there is an associated *importance value* which determines the significance of that question for a particular organization.

Table 4: Quantitative Input Survey Form

Quantitative Input Survey Form	
1.	Total existing workforce in each category.
2.	Percentage of each category of W/F lying in the nearing retirement age criteria Input: 1-10 (Scaled)
3.	What is the required male female ratio? Input: 1-10 (Scaled)
4.	Does the organization prefer to have employee from the closer regions for more dedicated workforce. Input: 1-10.
5.	Does the organization prefer multilingual employees? Input: 1-10.
6.	What is the perception of the management about their W/F Education Training level? Qualification Level Input: 1-10
7.	What is the expectancy level from the employees? Expectancy Input: 1-10
8.	How strong is the company policy for involuntary retirement of workforce? Input: 1-10
9.	Recruitment Process (Hiring the qualified W/F or development of skills within the organization?) Input: 1-10
10.	Do the employees find current job satisfactory according to the education/training level they have? Satisfaction Input: 1-10
11.	Are employees willing to switch over to another job because they find their current job status unsatisfactory as compared to their education level? Input: 1-10
12.	Is the organization providing training to the employees for enhancing their current job? Input: 1-10
13.	Is the organization providing multi-skilled training to the employees? Input: 1-10
14.	Is the organization expecting a growth in near future? Input: 1-10
15.	What is the percentage of employees leaving the company due to: Voluntary retirement Input: 1-10 (Scaled)
16.	What is the percentage of employees leaving the company due to: Involuntary retirement Input: 1-10 (Scaled)
17.	What is the percentage of employees leaving the company due to: Dissatisfaction Input: 1-10 (Scaled)
18.	What is the ratio of number of positions open for immediate hire vs. total positions available? Input: 1-10 (Scaled)
19.	The hiring activities planned by employers for the next six months Input: 1-10
20.	Management Level Question: Are the employees Over Paid/Under Paid? Input: 1-10
21.	Employee Level Question: Are the employees Over Paid/Under Paid? Input: 1-10
22.	Quantitative parameter : Input the number of the categorized W/F turnover rates in each category Input: 1-10 (Scaled)
23.	Categorized Quantitative data. (Percentage) for persons entering industry Input: 1-10 (Scaled)
24.	Categorized Quantitative data. (Percentage) for persons taking relevant training Input: 1-10 (Scaled)
25.	Categorized Quantitative data. (Percentage) completion rate of vocational training Input: 1-10 (Scaled)

26.	Ratio of Population growth Input: 1-10 (Scaled)
27.	Ratio of Migration (Employees switching between companies) Input: 1-10 (Scaled)
28.	Are the external job markets plentiful? Input: 1-10
29.	Quantitative Parameter: Current unemployment rate Input: 1-10 (Scaled)
30.	Ratio of Economic Growth Input: 1-10 (Scaled)
31.	The management is building organizational commitment by: Establishing an appropriate organizational culture Input: 1-10
32.	The management is building organizational commitment by: Improving opportunities for training, development Input: 1-10
33.	The management is building organizational commitment by: Introducing/encouraging employee participation Input: 1-10
34.	Is the management and leadership involved for improvement to safety and risk reduction Input: 1-10
35.	The management is promoting job satisfaction by: Remuneration Input: 1-10
36.	The management is promoting job satisfaction by: Working condition Input: 1-10
37.	The management is promoting job satisfaction by: Work organization Input: 1-10
38.	Are the management and leadership involved for improvement to safety and risk reduction? Input: 1-10
39.	How are employees recognized by the management? Verbal praise Input: 1-10
40.	How are employees recognized by the management? Salary increment Input: 1 to 10
41.	How are employees recognized by the management? Company benefits Input: 1-10
42.	Rate the following factor which influences organization's ability to attract potential employees. Information about the job Input: 1-10
43.	Rate the following factor which influences organization's ability to attract potential employees. Characteristic of the organization Input: 1-10
44.	Rate the following factor which influences organization's ability to attract potential employees. Subjective assessment about job and the organization Input: 1-10
45.	The management is changing employee perception of alternative employment opportunities by: Improving recruitment practices Input: 1-10
46.	The management is changing employee perception of alternative employment opportunities by: More effective communication and selective re-engagement of returners Input: 1-10
47.	Is the organization promoting any special programs? Input: 1-10
48.	Is the organization responding accordingly to the sudden changes in the W/F demand? Input: 1-10

49.	How is the relationship of the employees with the supervisor and coworkers? Relationship Input: 1-10
50.	The management is reducing the ease of movement to other jobs by: Financial participation Input: 1-10
51.	The management is reducing the ease of movement to other jobs by: Family-friendly employment policies; Work-life balance schemes Input: 1-10
52.	The management is reducing the ease of movement to other jobs by: Training infirm-specific skills Input: 1-10

10.2 Parameter Identification

Each of these questions has been carefully designed to describe the following 17 parameters as listed in Table 5. Answers to these questions should be searched in historical database of the organization to collect time-series data. Cross-sectional data can be collected by giving these questions to researchers, who keep an eye on current trends of similar types of other organizations. These easy-to-answer questions should be responded in three folds. First, what is the importance of a particular question for the organization (on the scale of 1-10); second, what is the actual answer of the question on same scale (1-10); and third, what was the total workforce size and mix corresponding to a one set of parameters (questions).

Table 5: Parameter List

1. Age	10. W/F Data 3: Completion rate of vocational Training
2. Gender Ratio	11. Population Growth
3. Geography	12. Migration
4. Educational Level	13. Size of industries in regional areas
5. Size of Business	14. Unemployment Rate
6. Remuneration	15. Economic Growth
7. Turnover rate	16. Attraction/Retention Parameter
8. W/F Data 1: Persons entering the industry	17. Emotional Demand of Work
9. W/F Data 2: Persons taking relevant training	

10.3 Fuzzy Logic Controller

Application of fuzzy set theory to represent non-statistical uncertainty and approximate reasoning in real life situations can be found in literature (Gen et al., 1996; Mierswa, 2005). The fuzzy logic controller (FLC) is a conceptual input-output machine, capable of making certain decisions based on the input linguistic variables fed to it. In this research, the FLC is structured as Multiple-Input-Single-Output (MISO) for making relevant decisions based upon the two input linguistic variables and the rule-base associated with them.

After formulation of the rule base, the next step is to determine the membership functions of the input and output parameters. In general, each membership function is defined by a start and an end point and can take either linear or non linear form, in which Triangular, Gaussian and Sigmoid functions are the most prevalent (Ross, 1997). In this research, the inputs are fuzzified (Ross, 1997) using the Gaussian membership functions distributed symmetrically across the universe of discourse. Three fuzzy sets are used in each case with an overlapping of about 25% - 50%. If-Then rules (Figure 12), which are derived from the rule base, are used to compute the fuzzified output that is ready for defuzzification. In this research, we have used the weighted average mean methodology for defuzzification.

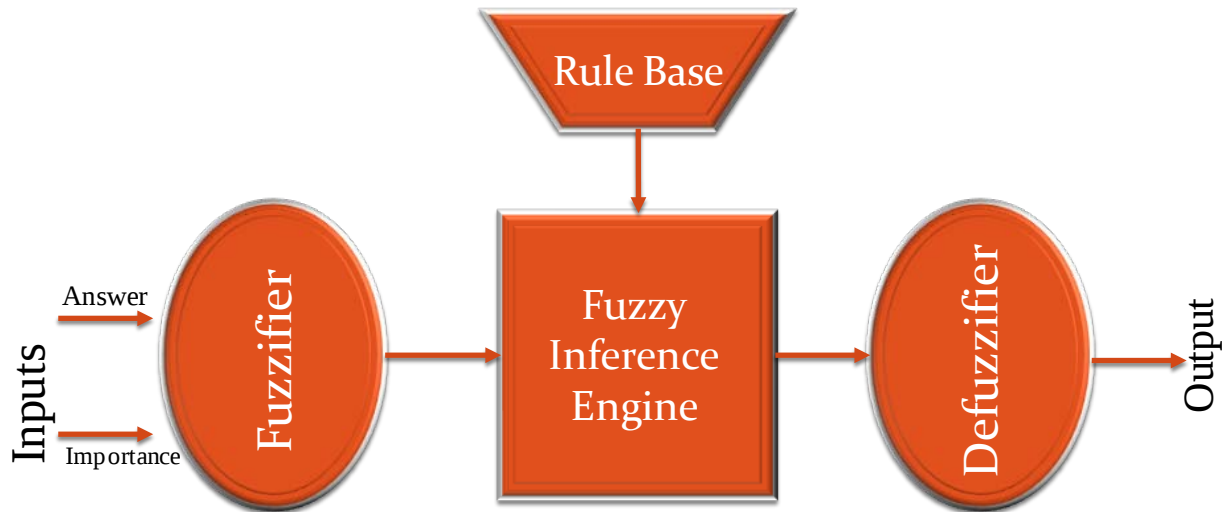


Figure 12: The Fuzzy Logic Controller

A fuzzy logic controller (as shown in Figure 12) is utilized to amalgamate the information from the importance value and the actual answers to get partial quantified parameter value. If two or more questions contribute towards one parameter then the average value of all the partial quantified values are utilized to determine the quantified parameter value.

One complete set of information consists of quantified values of the 17 parameters and the corresponding workforce size and mix for one year (or one month). Several similar sets of such information are collected in the form of time series data and cross-sectional data so that some meaningful information can be derived from it.

After we have this complete quantified dataset, the data collection and parameter identification phase is over and we move to the forecasting phase of data mining algorithm.

11.0 SOFT COMPUTING AND ARTIFICIAL INTELLIGENCE

Over past decades there has been a dramatic increase in the interest and use of various soft computing techniques for scientific and engineering techniques pertaining to artificial intelligence. In general, a system is defined as an artificially intelligent system (or intelligent system in short) (Sanchez et al., 1997), if it covers approaches to:

- Design,
- Optimization,
- Control and coordination of various sub-systems working in synergy without requiring any specific mathematical model.

The concept of intelligent systems is driven by the evolutionary environmental phenomena around us. It works in a way similar to how different living organisms survive in the environment. In essence, such systems derive concepts from evolution of human beings, simulation of neural network within the human brain, food foraging behavior of insects, immune system within the human beings and DNA patterns etc. Besides, simulation of the actual cooling/annealing process and other physical laws in nature also help us in solving complex problems and are, therefore, grouped in the artificially intelligent techniques.

A number of research contributions can be found in literature addressing various disciplines of these artificially intelligent soft computing techniques (Anderson and Ferris, 1989; Quagliarella, 1998; Dorigo, 2004; 2006; 2008; Dorigo and Stützle, 2004). However, most of them focus only on certain areas of soft computing techniques and applications. In this research, an effective intelligent decision support system is constructed by combining several soft computing techniques such that they work in synergy to predict the future workforce for an organization.

This research is dedicated to providing the highlights in current research in the theory and applications of artificial intelligence techniques in prediction and forecasting scenario. This section begins with the introduction of various artificial intelligence techniques including artificial neural network, decision trees, fuzzy logic and several evolutionary computing techniques.

11.1 Artificial Neural Networks

In recent years, a great deal of attention has been focused on the idea of predicting and analyzing the hidden pattern within large datasets by utilizing Artificial Neural Networks (ANNs). ANNs can be realized as circuits of highly interconnected units with modifiable interconnection weights which mimic the learning capabilities of the human brain. In general they represent an input output machine, with a connection of simple processing elements capable of processing information and recognizing the hidden pattern within the data.

11.1.1 Working of Neurons within a Human Brain

In general, neural network technique is adapted to build an intelligent program (to implement intelligence) using models that simulate the working network of the neurons in the human brain (Hopfield, 1982; Hopfield and Tank, 1985). A generic diagram of a neuron in a real neural

network is shown in Figure 13. There are several parts of neuron that need description. Several protrusions coming out of neuron are called dendrites and a long central branch is known as the axon. A neuron is joined to other neurons through the dendrites. The dendrites of different neurons meet to form synapses, the areas where messages pass. The communication within the neurons is setup by receiving the impulses via the synapses.

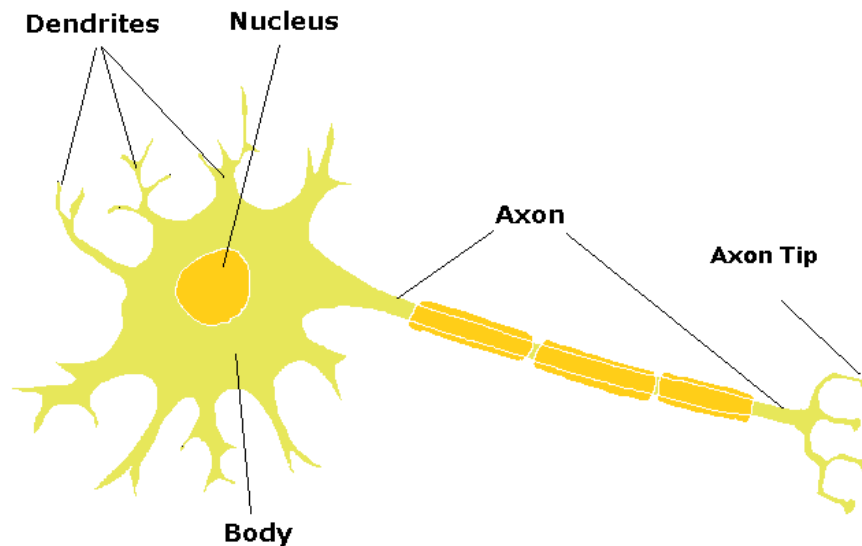


Figure 13: Structure of Neuron

The working of the neural network is governed by the impulses transferred. It operates by simple if then rules within the decision making. To be specific, if the total of the impulses received exceeds a certain threshold value, the neuron sends an impulse down the axon. This axon is connected to other neurons through more synapses and thus the message is passed on. The synapses work twofold and may have two different natures: excitatory or inhibitory. An *excitatory synapse* acts as an additive synapse in a sense that it adds to the total of the impulses reaching the neuron. Similarly, an *inhibitory synapse* reduces the total of the impulses reaching the neuron. In a global sense, a neuron receives a set of input pulses and sends out another pulse that is a function of the input pulses.

11.1.2 Simulating Neural Networks for Computation and Learning

The working behavior of neurons in the human brain can be utilized in performing complex computation jobs with the help of developing processing power of the computers. As discussed earlier, the basic building block of any neural network (either human brain or ANN) is a neuron. According to Russell and Norvig (1995) a neuron is defined as a cell in the brain whose principal function is the collection, processing, and dissemination of electrical signals. It has been proven that the neurons are responsible for the human capacity to learn. Therefore, the physical structure of a neuron is being emulated by an artificial neural network to accomplish machine learning or data mining in broader perspectives.

In ANNs, similar to biological neurons, artificial neurons are connected to several other neurons forming massive interconnections within them. These artificial neurons are capable of accepting information from other artificial neurons and also transmitting information to other neurons. The neurons within human brain work by transmitting electrical signals, whereas, artificial neurons receive a number from other neurons, and process these numbers accordingly.

Each artificial neuron has fundamentally the same structure and they look alike. Neurons can perform simple calculations based on the inputs they receive and the inbuilt functions associated with them. However a unique and most powerful feature about these artificial neurons is that they can solve computationally complex problems by their ability to perform calculations collectively. Thus, by combining groups of neurons, one can perform extremely complex operations, even while each neuron acts independently of others. It is for this reason that neural networks have often been described as “parallel distributed process” (PDP) networks.

A characteristic feature of these networks is that these networks are trained to solve problems by modifying how they receive information from other neurons that rather than being set up to solve a specific problem. Thus they act as a “black box” kind of function between input and outputs, where a black box represents a complex interrelated structure of neurons. The user need not know the structure of the trained neural network and can still operate the network easily. However, at the same time, one of the major drawbacks of such a network is their same black box quality: after a network is trained, it may be extremely difficult to interpret its state to learn which inputs and qualities are important to the solution received or to make a mathematical/statistical model out of it.

These neurons are connected to each other in a mesh like fashion forming multiple interconnections within the neural network. A network consists of several layers which are made up of neurodes or cells. To be specific, there are three distinct layers in the neural network: input layer, hidden layer(s) and output layer. Input neurons are ones that are fed data from an external source, such as a file, rather than other neurons. Hidden neurons accept their information from other neurons' outputs and pass it on to the next neuron. Output neurons are like hidden neurons, but rather than passing the processed information to a new set of neurons, they save the information and allow it to be read and interpreted by an external actor. The interconnections existing within the layers represent the information exchange within the layers and cells.

A network must include a learning model to equip it with the pattern recognition capabilities. In general, a neural network learns by evaluating changes in input. The general structure of NN representing the input layer, hidden layers and output layer is provided in Figure 14.

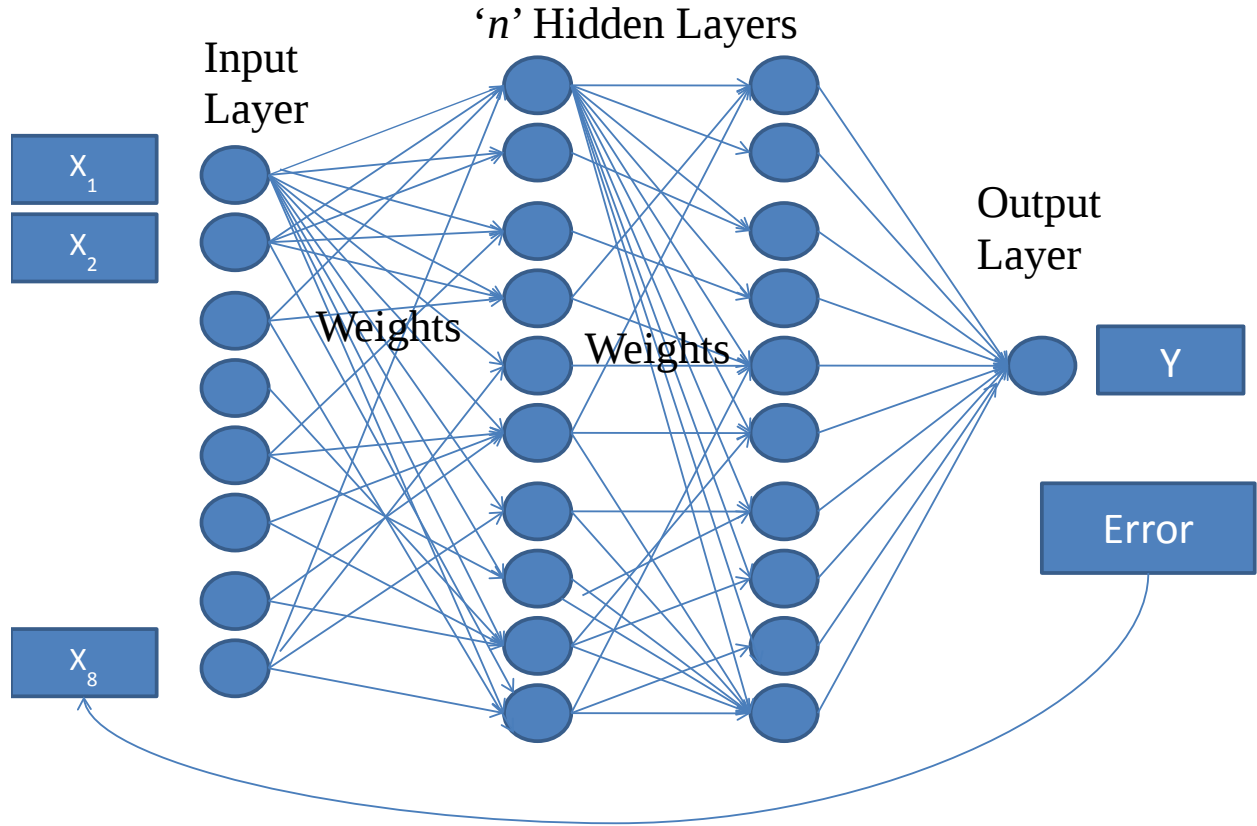


Figure 14: Neural Network Structure (Tripathi et al., 2010)

In this figure, there are 8 input neurodes within the input layer for 8 inputs in the neural network. There are 2 hidden layers and with 10 cells in each layer and again one in output layer.

We assume that we have T input patterns X_t available with us and corresponding T output patterns Y_t which are used to train the network. Thus, we have following relationship:

$$1 \leq t \leq T \quad (1)$$

The total error value is calculated by using a quadratic error function defined as below:

$$E = \sum_{t=1}^T (Y_t - F(X_t))^2 \quad (2)$$

where, F is the current function implemented by the network. This function F depends upon the current weights of the network. These weights are initialized either randomly or by certain criteria. During the training phase, a learning algorithm is utilized whose main objective is to reduce the quadratic error value (E).

Figure 15 illustrates an analog of a neuron as a threshold element. The variables $x_1, x_2, \dots, x_i, \dots, x_n$ are the n inputs to the threshold element. These inputs are analogous to impulses that arrive from several different neurons to this neuron. The variables $w_1, w_2, \dots, w_i, \dots, w_n$ are the weights associated with the impulses/inputs. These weights signify the relative importance that is associated with the path from which the input is coming. As discussed before, synapses can be excitatory and inhibitory. When w_i is positive, input x_i acts as an excitatory signal for the element. When w_i is negative, input x_i acts as an inhibitory signal for the element.

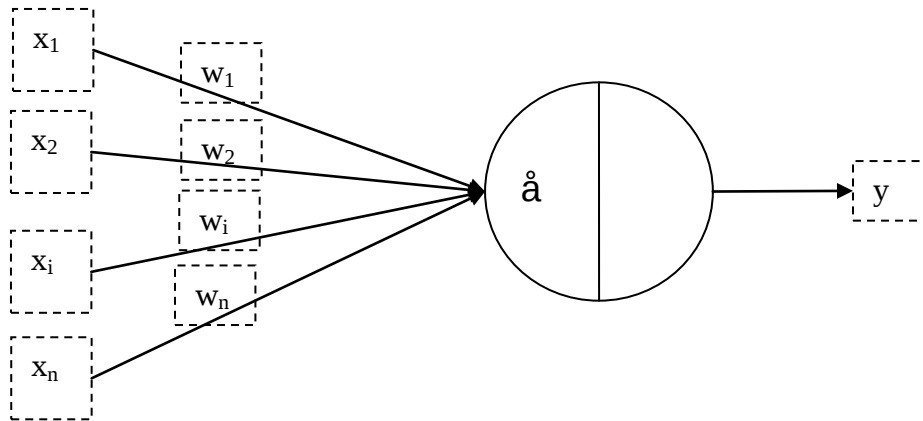


Figure 15: A Threshold Element as an Analog to the Neuron

The threshold element computes the weighted sum of these inputs as $\sum w_i x_i$ and if it surpasses the prescribed threshold value 'Th', an output using a non-linear function as shown below:

$$y = F \left(\sum w_i x_i - Th \right) \quad (3)$$

This non linear function $F(.)$ is a modeling choice by the user. A widely used function is sigmoid function as shown below:

$$F(z) = \frac{1}{1 + e^{-z}} \quad (4)$$

Other popular choices for $F(.)$ are step function, ramp function on the unit interval.

11.2 Genetic Algorithm

Genetic Algorithm (GA) is a powerful optimization tool that imitates the natural process of evolution and Darwin's principle of 'Survival of the Fittest' (Goldberg et al., 1989). Some of the typical applications of GAs includes: Traveling salesman problem (Grefenstette et al., 1985); Scheduling problem (Davis, 1985; Cleveland and Smith, 1989); VLSI Circuit layout design problem (Fourman, 1985); Computer aided gas pipeline operation problem (Goldberg, 1987a, 1987b); Communication network control problem (Cox et al., 1991); Real time control problem in manufacturing systems (Lee et al., 1997) etc. The efficient implementation of GA over a problem requires following:

- A genetic representation of solution which can be expressed as a string (chromosome);
- A way to generate an initial population of solutions.
- An evaluation function to compute fitness value corresponding to each string in the population.
- Genetic operators which can alter the genetic composition of the children solutions. There are three main operators in a genetic algorithm: selection, crossover and mutation.
- Adequate (tuned) values for various parameters of genetic algorithm.

Once the solution is encoded as a string and a measure of fitness evaluation is chosen, the GA proceeds as shown in the pseudo code (Figure 16).

```
Genetic Algorithm
{
    Initialization of population of chromosomes
    Evaluation of chromosomes
    While termination criteria is not satisfied
    {
        Selection of parent chromosomes to form mating pool
        Crossover operation to form offsprings
        Mutation in offsprings
        Evaluation of offsprings generated
    }
}
```

Figure 16: Pseudo Code of Genetic Algorithm

A few important terminologies utilized in GA are discussed as below.

11.2.1 String (Chromosome) Representation

The first step in the implementation of a search technique to any problem is the representation of search space in terms of algorithmic parameters. The string encoding can be:

- Binary
- Real Number
- Integer or literal permutation encoding
- General data structure encoding.

According to the structure of the problem and requirements, the encoding structure can also be classified into following two types:

- One dimensional encoding: This is the most common form of encoding and is also known as the linear array encoding.
- Multi dimensional encoding: In more complex real world problems, the encoding is done in form of matrix or even more dimensions.

11.2.2 Selection

The basic idea behind the genetic algorithm is natural selection and survival of the fittest. With a directed search, the search will end prematurely whereas for a more randomized search, it will take too much time for the evolutionary algorithm to search. The common types of selection procedures (Deb, 2004) are:

- Roulette wheel selection: This is the best known selection procedure where the underlying idea is to assign the selection probability or the survival probability for each chromosome proportional to the fitness value.

- $(m+1)$ -selection: As opposed to the roulette wheel selection, this selection procedure selects a deterministic procedure that selects the best chromosome from the current and children generations.
- Tournament selection: First of all a tournament size is selected and then individuals are randomly assigned to the tournament from the current and the children generation. Thereafter, the best individual from these is said to win the tournament and is retained.
- Ranking: In this selection, both- the current and the children - generations are sorted in the decreasing order of their fitness value. Thereafter, the top population size individuals are selected and passed on to the next generation.
- Elitist strategy: In this strategy, the best individual never dies and is always retained. It can be mixed with all four selection procedures listed above.

11.2.3 Crossover

In crossover, two parent strings combine to create new, hopefully better, offspring. A number of crossover methods have been designed (Goldberg 1989; Spears 1997) to achieve the fusion of genetic materials like, k-point crossover, uniform crossover, uniform order based crossover, order based crossover, partially matched crossover (PMX), cycle crossover (CX), precedence preservative crossover (PPX) etc. In PPX, a mask is first created that consists of random 1s and 2s corresponding to parent part from which the information has to be taken. For example, if the mask for child 1 reads 22221211, then the first four genes of second parent would make up the first four genes of child 1; the fifth gene of child 1 will comprise that first gene of parent one which is not present in child 1. The process is depicted in Figure 17.

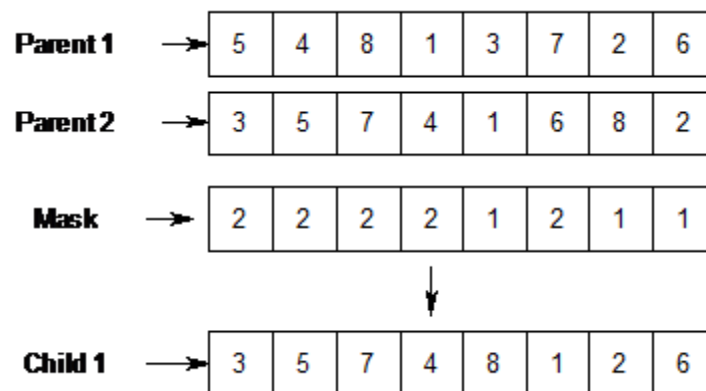


Figure 17: PPX Example

11.2.4 Mutation

Mutation is mainly used to impart exploitation in genetic algorithms and thereby prevent them from premature convergence. It involves random changes in genes comprising the chromosome. Swapping mutation has been utilized for the problem at hand which randomly selects two bits and then swaps them.

In general, a genetic algorithm always maintains a fixed population of individuals in each generation. Each individual within this population represent a potential (feasible) solution (either good or bad) to the problem at hand. Thereafter each individual is evaluated for its fitness using

the fitness function. The fitness function is designed in such a way that for maximization and minimization problems, the maximum fitness function values are preferred. This is called the Darwin's 'survival of the fittest' theory. Thereafter, by using reproduction, new children (individuals) are produced and only the ones with higher fitness values are allowed to live by some selection process. The process is iterated for several generations (several hundred or thousand generations) before it is stopped. The stopping criteria can be set as maximum number of iterations/functional evaluations. Most of the time, the stopping criterion is based upon the user satisfaction level and experience.

11.3 Algorithm of Self Guided Ants

Considering the computational complexity involved in various problem types, this research utilizes a new meta-heuristic termed as Algorithm of Self Guided Ants (ASGA) which utilizes a speed versus accuracy tradeoff in collective decision making demonstrated by *Temnothorax Albigipennis* ants (Franks and Richardson, 2006). Here we mathematically model and simulate the house hunting behavior of these ants and map it with the selection of the edges joining two nodes by utilizing a new state transition rule. An interesting phenomenon occurring in colonies of such ants corresponds to their quick decision making in harsh environmental conditions than benign ones (Dorigo and Gambardella, 1997). However, it was observed that the quick decision making generally leads the ants to an erroneous choice (i.e., making choice of comparatively weaker nest). Recent researches over *T. Albigipennis* ant have empirically established that this error in decision making is primarily due to the judgment error and not due to the omission error as ants have discovered almost all alternative nests before making a choice for initiating an emigration (Franks et al., 2003). Thus, this error in accuracy is mostly due to the strict time constraints imposed during the need for urgent decision making in harsh environmental conditions. By incorporating similar idea in foraging behavior of ants, convergence trends of the ASGA can be controlled making an opposite tradeoff in solution quality.

The choice of an appropriate nest is a highly structured process involving following stages; exploration, decision making and migration. These stages are characterized by multifaceted options, high stakes and involvement of numerous individuals (Franks et al. 2003). The crisis management and rapid achievement of consensus over optimal location of nest from amongst a large set of probabilities are simultaneously taken care of by these ants. For this purpose, the ants utilize the concept of quorum threshold (detailed in Section 11.3.2) while choosing whether to give emphasis to speed or to move for accuracy.

Moreover, to simulate the ant colonies more realistically, this research presents an idea of self guided ants and utilizes an innovative and effective strategy for selecting the next node in the ant's tour. The underlying idea behind this selection strategy resides on the fact that while making a transition to the next node, the ants take into account only the pheromone concentration. Therefore, in ASGA, the consideration for distance between the two nodes is omitted from the criteria of the probabilistic selection of the nodes. However, this heuristic information has been implicitly merged with pheromone concentration by utilizing some appropriate dynamic pheromone evaporation rule during the tour construction. Based upon this, the ASGA is first discussed in terms of a traveling salesman problem (TSP) and thereafter its application on the FDSS problem is detailed. Mathematically, structure the ASGA is as follows,

11.3.1 The Initialization Criteria

This research assumes the path between any two nodes as an arc with an integrated collection of infinite discrete points. Simulating the actual environment of ant movement, the initial pheromone count at each arc is assumed to be equal. Therefore, the pheromone concentration on the path between the two nodes is calculated by computing the ratio of the net pheromone content on the arcs and the length of that arc. If $l(r, s)$ denotes the length of the path between the nodes r and s then the initial pheromone concentration $\tau_0(r, s)$ between the same two paths is modeled by the following equation,

$$\tau_0(r, s) = \frac{Q_0}{l(r, s)} \quad (5)$$

In Equation 5, Q_0 is a fixed pheromone content assumed to be distributed uniformly over each arc. Computationally, it is a model dependent parameter. For initialization, ants are randomly positioned on the feasible nodes. In order to simulate real ant's motion in forward direction only, a data structure called tabu-list is associated with each ant for maintaining the feasibility constraint.

11.3.2 The Self Guided State Transition Rule

In the proposed ASGA, the artificial ants utilize a concept similar to quorum threshold (Franks et al., 2003) as shown by the *T. Albipennis* ants. For this purpose, a quorum index is defined whose values determine whether to use probabilistic model to select the next nodes in the path or to simply follow the pheromone trail intensity. If there is an urgent need for speedy convergence (harsh environment), the quorum threshold is low, otherwise it is high (benign environment). Therefore, before making a move, the ant decides its behavior for running – *tandem running* or *social carrying*.

The scout first assesses the quorum size (or the quorum index), which is compared with the context sensitive threshold quorum ' B '. If the threshold is low, tandem running is preferred, whereas social running is used in the other case. The probabilistic model simulates the tandem running whereas the mechanism to follow pheromone trail intensity is used for social carrying. The quorum index for an ant positioned at node r is defined as,

$$b_{rv} = \frac{N_{rv}}{N} \quad (6)$$

where, N_{rv} is the number of ants that have passed through the nodes (r, v) in the last iteration; N is the total number of ants in the colony. The quorum index is calculated for each path connecting the current node r and the feasible node v .

Mathematically, the state transition rule for the path between the nodes r and s is dynamically chosen by the ant according to,

$$\text{State transition rule}(r, s) = \begin{cases} \text{Tandem Running rule} & \text{if } b_{\max}(r, v) > B \\ \text{Social Carrying rule} & \text{otherwise} \end{cases} \quad (7)$$

where, $b_{\max}(r, v)$ is the maximum quorum index for the ant at node r . Whenever an ant arrives at any node r , the decision regarding choosing the next node is based upon the maximum quorum

index and the value of quorum threshold B . If the value of maximum quorum index is larger than the quorum threshold, tandem running is invoked, otherwise social running is used. These two state transition rules are defined below.

Tandem Running rule:

In any iteration y , let T be the set of cities which satisfy the criteria for quorum threshold ($b(r,v) > B$). This rule is followed if $n(T) > 0$ ($n(T)$ represents the cardinality of set T), i.e., at least one $b(=b_{max})$ exists such that $b_{max}(r,v) > B$. In such case the transition probability is given as,

$$p_y(r,v) = \begin{cases} \frac{b_y(r,s)}{\sum_{t \in T} b_y(r,t)} & t \in T \\ 0 & \text{otherwise} \end{cases} \quad (8)$$

In tandem running rule, the speed accuracy tradeoff is maintained where the transition probabilities governing the movement of ants derives information only from the quorum index and the quorum threshold value. Such a modification maintains fast convergence of the algorithm in addition to keeping the explorative structure in cases of local convergence.

Social Carrying rule:

This rule is followed if $n(T) = 0$, i.e., if no b exists for which the criteria for quorum threshold ($b(r,v) > B$) is satisfied. In this case, the transition probability is mathematically modeled as,

$$p_y(r,v) = \begin{cases} \frac{\tau_y(r,v)}{\sum_{u \in J_y(r)} \tau_y(r,u)} & \text{if } v \in J_y(r) \\ 0 & \text{otherwise} \end{cases} \quad (9)$$

Thus, in any iteration y , an ant positioned on node r chooses to move to a feasible node v (i.e., a node lying in the list of feasible nodes $J_y(r)$), according to eq. (9).

11.3.3 The Pheromone Update Rules

Two pheromone update rules are utilized to update the pheromone concentration lying between the paths of the nodes.

Dynamic Pheromone Evaporation and Update rule:

Immediately, after the choice of the next node has been made by an ant, the content of its tabu-list is updated and a local pheromone update rule is applied. First the pheromone level of all the paths is allowed to evaporate equally.

$$\tau^{new}(i,j) = (1 - \rho_l) \cdot \tau^{old}(i,j) \quad (10)$$

where, ρ_l is the local pheromone evaporation parameter. Now, additional pheromone is given to the paths over which each ant has passed by. The study of real ant phenomenon states that the ants reside for lesser time on the shorter paths than on the longer paths. Thus, evaporation on the shorter paths occurs for a lesser time which indirectly leads to higher pheromone concentration

over these paths than on the longer ones. To simulate this phenomenon, cumulative pheromone added to each path is governed by the following equation,

$$\tau^{new}(r,s) = \tau^{old}(r,s) + \rho_l \cdot \left(\frac{Q_l}{l(r,s)} \right) \quad (11)$$

where, Q_l is the model dependent positive quantity. This dynamic pheromone evaporation and update rule implicitly merges the greedy heuristic information with the pheromone information by logically relating it with the actual evaporation phenomenon that occurs in movement of real ants.

Global Pheromone Update rule:

Moreover, a global update rule is followed at the tour completion by all the ants to reward the ant which corresponds to the best tour till the current iteration. The purpose of this rule is to search in the vicinity of the best tour found so far.

$$\tau^{new}(i,j) = \tau^{old}(i,j) + \rho \Delta \tau(i,j) \quad (12)$$

$$\Delta \tau(i,j) = \sum_{k=1}^N \Delta \tau^k(i,j) \quad (13)$$

$$\Delta \tau^k(i,j) = \begin{cases} Q \cdot \left(\frac{1}{L_k \cdot l(i,j)} \right) & \text{if } (i,j) \in k^{th} \text{ ant's tour} \\ 0 & \text{otherwise} \end{cases} \quad (14)$$

Parameter Q is the positive model dependent constant and L_k is the tour length by the best ant. After global pheromone update, all the ants are reinitialized and their tabu-lists are emptied. Algorithm is iterated in this fashion till certain terminating criterion (maximum number of iterations or functional evaluations) is met.

11.4 Self Guided Ant-based Genetically-Optimized-Neural-Network (SGA GONN) Algorithm for DFSS

As discussed earlier, developing a generic analytical/mathematical model for workforce planning using an objective function and a set of constraints is not feasible and leads to a model inconsistent with several other organizations. Therefore, an artificial intelligence based decision support system has been proposed in this research. In general, artificial neural networks have been proved to be universal approximators and a three-layer feedforward neural network can approximate any nonlinear continuous function to an arbitrary accuracy (Brown and Harris, 1994). Owing to its structure, a neural network utilizes several learning algorithms, e.g., GA (Tripathi et al., 2008) and backpropagation (Pham and Karaboga, 2000) for enhancing its prediction capabilities. This research utilizes GA as learning algorithm and enhances the learning capability of the network by incorporating the concept of self guided ants in the learning procedure. The structure and learning steps of the proposed SGA GONN are as follows:

- 1) First, a network structure is defined with a fixed number of inputs, hidden nodes and outputs. The network is initialized randomly with weights w_i for each connection i in the network.

- 2) Second, GA is chosen to realize the learning process such that each connection in the network has its own mutation rate. However, a fixed structure may not provide the optimal performance within a given training period as the network may have some of its connections redundant and some other connections highly significant (Yao, 1999).
- 3) To obtain the connections which significantly contribute to the network performance, artificial agents called self guided ants are stationed on the input nodes and initial pheromone level $t_i = t_0$ is maintained on the paths connecting the two nodes. It should be noted that net weight for any connection is determined by the expression $t_i' \cdot w_i$ and not w_i .
- 4) After each generation of GA, the self guided ants traverse through the network (according to the state transition rule as shown in eq. 7) and evaluate the change in the weight of each connection ($\Delta t_i' \cdot w_i$) and its corresponding effect on the network performance. If change contributes towards increasing network performance then the mutation rate for that connection is increased.

$$p_y(r,v) = \begin{cases} \frac{\tau_y(r,v)}{\sum_{u \in J_y(r)} \tau_y(r,u)} & \text{if } v \in J_y(r) \\ 0 & \text{otherwise} \end{cases} \quad (15)$$

where, $p_y(r,v)$ is the transition probability of an ant stationed on node r to move to node v and $J_y(r)$ is the tabu-list maintained to ensure that ants do not travel a node twice within a tour.

- 5) Pheromone update rules similar to are utilized to evaporate/update the pheromone concentrations on the connections between the nodes.

$$\tau_i^{new} = (1 - \rho) \cdot \tau_i^{old} \quad (16)$$

$$\tau_i^{new} = \tau_i^{old} + \rho \cdot \left(\frac{Q_l}{e} \times \Delta(\tau_i \times w_i) \right) \quad (17)$$

where, ρ is the pheromone evaporation parameter, Q_l is the dynamic model dependent parameter and e is the error value from the network. Therefore, the connections which do not contribute significantly towards the network performance die slowly over the period of time due to reduced pheromone concentrations. Figure 18 shows structure of SGA GONN with 2 inputs and two outputs.

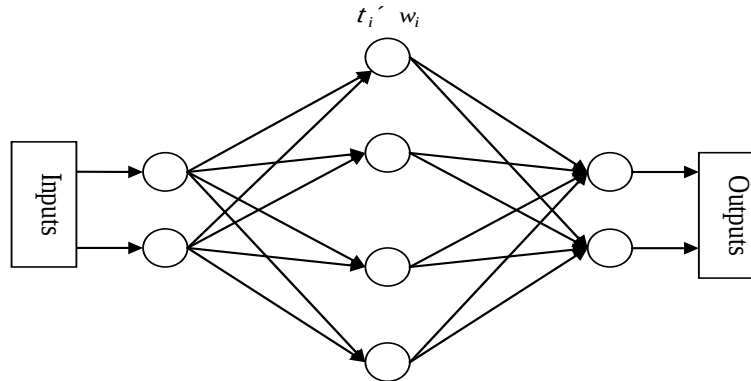


Figure 18: SGA GONN with Weights t_i' and w_i

11.5 Clonal C-Fuzzy Decision Tree (C²FDT)

The architecture of the decision trees involves a comprehensive list of various training and pruning schemes, a diversity of discretization algorithms, and a series of learning refinements (Dobra and Gehrke, 2002; Weber, 1992; Yildiz and Alpaydin, 2001). Beside the good capabilities of classification and discrete prediction, decision trees are allied with the following shortcomings and limitations. First, deficiency refers to its capability of operating on only discrete attributes having a finite number of values (Pedrycz and Sosnowski, 2000). Second, in the design process, the most discriminative attribute is selected at a time and the tree is expended by adding the node whose attribute's values are located at the branches emerging from this node (Pedrycz and Sosnowski, 2005). Moreover, in expansion phase, consideration of one attribute at a time also gives rise to problem of fragmentation, repetition and replication resulting in complex and large decision tree. Third, generally decision trees are suitable only for discrete class problems, while continuous prediction problems are handled by regression trees (Han and Kamber, 2001). Discretization plays a very important role in handling of continuous attributes and it directly impacts performance of the tree. In this sense, the way in which tree has been grown and the sequence of attributes selected are inherently affected by the discretization mechanism. Clearly, these two design steps cannot be disjointed.

In order to overcome these deficiencies, researchers like Eggermont (2002), Mendes et al., (2001) have developed fuzzy decision trees with genetic programming, clustering, and fuzzy classification rules. Pedrycz and Sosnowski (2005) have developed cluster based fuzzy decision trees (C-fuzzy decision trees). In the development of cluster based fuzzy decision trees Fuzzy C-Mean Clustering (FCM) were treated as its generic building blocks. FCM is an iterative optimization process for information granulation, where partition matrix and a prototype are regularly updated until some termination criterion has been met. This algorithm has been exhaustively elaborated by Bezdek (1981) and Pedrycz and Sosnowski (2000). The FCM uses calculus based optimization methods; therefore it may be trapped in local extrema during optimizing the clustering criterion. It is also very sensitive towards the initialization. Due to these deficiencies of calculus based optimization researchers like Tseng and Yang (2001) and Maulik and Bandyopadhyay (2003) have made use of evolutionary algorithm to optimize these prototypes. For optimizing the same prototype Shukla and Tiwari (2009) used genetic algorithm.

One of the aims of this research is to develop a class of new decision trees “Clonal C-Fuzzy Decision Trees (C²FDT)” which are able to ameliorate the deficiencies allied with existing architectures. Similar to as Shukla and Tiwari (2009), in this research, information granule have been considered as one of the fundamental building blocks of the C²FDT. The real difference between C-fuzzy decision trees and C²FDT lies in encompassing the clustering methodology. In contrast to C-fuzzy decision trees where only FCM acts as generic building block, this research uses clonally optimized fuzzy clustering for the construction of the tree. The proposed architecture attempts to achieve both the avoidance of local extrema and minimal sensitivity to initialization. FCM algorithm and clonal algorithm is discussed in next subsections.

11.5.1 Fuzzy C-Means (FCM) Clustering

FCM is an information granulation technique, and in the past it is comprehensively described by Pedrycz and Sosnowski (2000) and Bezdek (1981). However, for the sake of convenience, we are presenting here a brief idea of this well known algorithm.

The FCM algorithm is an example of an object oriented fuzzy clustering where clusters are built through the minimization of some objective function $f(U, V)$. Let $X = \{x_1, x_2, \dots, x_N\}$ be the set of categorical objects, where N is the number of data points and let each set of data is consist of S attributes $A = \{A_{11}, A_{12}, \dots, A_{1S}\}$. Let, $V = \{v_1, v_2, \dots, v_c\}$ is a set of corresponding clusters center in the data set X , where c is the number of clusters.

$$f(U, V) = \sum_{i=1}^c \sum_{k=1}^N \mu_{ik} \left(\|x_i - v_k\| \right)^2 \quad (18)$$

$f(U, V)$ is a squared error clustering criterion, and solutions of equation (18) are least squared error stationary points of $f(U, V)$. μ_{ik} is the membership grade of data x_i to the k^{th} cluster.

Where, $U = (\mu_{ik}), i = 1, 2, \dots, c, k = 1, 2, \dots, c$ is a fuzzy partition matrix, $\|x_i - v_k\|$ represents the Euclidean distance between x_i and v_k . Meanwhile, μ_{ik} has to satisfy the following conditions:

$$\mu_{ik} \in [0, 1], \forall i = 1, \dots, n, \forall j = 1, \dots, c \quad (19)$$

$$\sum_{j=1}^c \mu_{ij} = 1, \forall i = 1, \dots, N \text{ and } k = 1, 2, \dots, c \quad (20)$$

Prototype of the cluster is treated as a typical vector or representative of data forming it. These data come as ordered pairs $\{x(k), y(k)\}$.

In brief, the traditional FCM is an iterative optimization process in which the partition matrix and prototypes are iteratively updated. Generally, U and V are governed by following expression Bezdek (1981):

$$\mu_{ik} = \frac{1}{\sum_{j=1}^c \left(\frac{\|x_i - v_k\|}{\|x_j - v_k\|} \right)^{2/(m-1)}} \quad (21)$$

$$v_i = \frac{\sum_{k=1}^n \mu_{ik}^m \cdot z_k}{\sum_{k=1}^N \mu_{ik}^m} \quad (22)$$

Parameter m is the fuzzification index and is used to control the fuzziness of membership of each datum in the range $m \in [1, \infty]$. $m = 2$ is chosen for all the experimental purposes. Although there is no theoretical basis for the optimal selection of m , but in the literature same is used by the researchers. The iterative process starts from a randomly initiated partition matrix and involves the calculation of the partition matrix and prototypes.

11.5.2 Clonal Algorithm

In last few decades the researchers, involved in designing and optimization of engineering systems, have paid considerable attention to the evolutionary processes like Artificial Immune System (AIS), Genetic Algorithms (GA), DNA Computing, Ant-Colony Optimization (ACO). AIS provides a way to deal with the complex computational problems like pattern recognition, elimination, machine learning and optimization. Clonal Algorithm is devised by principals of artificial immune system (De Castro and Timmis, 2002). Therefore, it is essential to first discuss the preliminaries of the immune system.

The vertebrate immune system is a complex system having large number of functional components. The main task of the immune system is to survey the organism in search of malfunctioning cells from their own body and foreign substances that are recognized by system called antigens. The constituents of the immune system that recognize antigens are called antibodies. There are two major categories of immune cells: B cells and T cells. B cell recognizes antigens free in solution (in blood stream) on the other hand T cells require antigens to be presented by accessory cells. After generation of T-cells they migrate in to thymus where they mature. During maturation, all T-cells that recognize self-antigen are excluded from population of T-cells this process is termed as negative selection (Nossal, 1994). If a B-cell interacts with a non-self antigen with affinity threshold; it proliferates and differentiates in the memory and effector cells, this process is termed as clonal selection (Ada and Nossal, 1987). In contrast, if a B-cell recognizes a self-antigen, it might results in suppression as recommended by the immune network theory (Jerne, 1974).

When a non-self antigen is recognized by a B-cell receptor with threshold affinity, it is selected to proliferate and produce high volume of antibodies as shown in Figure 19. During reproduction the B-cell progenies (clones) with strong selective pressure participates in hypermutation process. The whole process of mutation and selection is known as maturation of the immune response. Depending upon affinity, a selection is made from matured pool of antibodies. This results a high quality solution that exhibit prudence. From an engineering point of view this is the most alluring characteristic of the immune system because the candidate solutions with higher affinity must somehow be preserved as high quality candidate solutions and will only be replaced by matured clones. In a T cell dependent immune response, the repertoire of antigen-activated B cells is diversified by two mechanisms: hypermutation and receptor editing as shown in Figure 20.

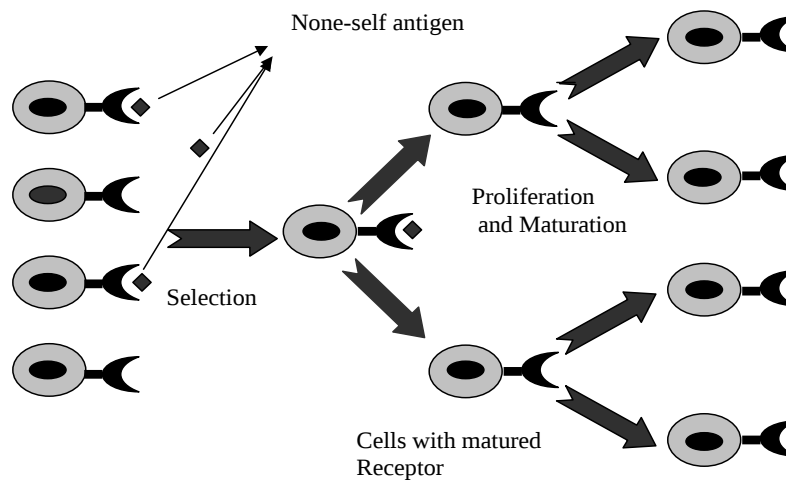


Figure 19: Schematic Representation of Proliferation and Maturation (Shukla et al., 2009)

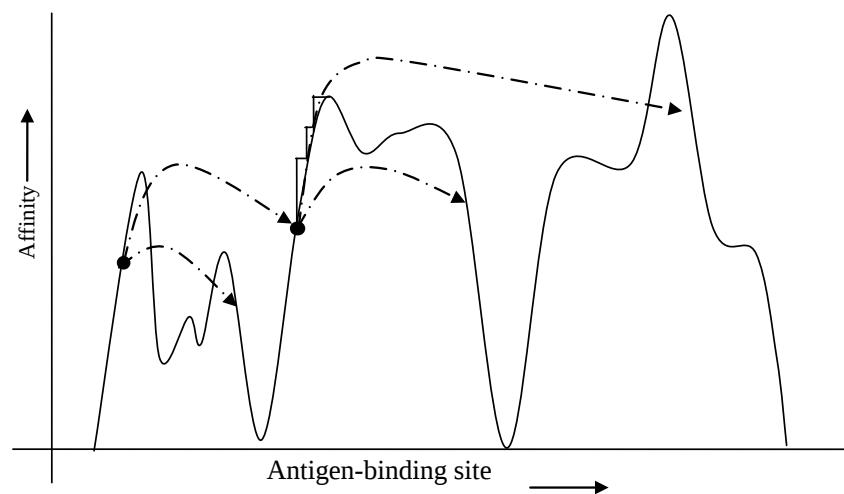


Figure 20: Antigen-binding Sites (Shukla et al., 2009)

Random changes are introduced in to the genes responsible for the Antigen-Antibody interaction, and occasionally one such change will lead to an increase in affinity of the antibody. The hypermutation operator works in a similar fashion to mutation (Shukla et al., 2009). The difference lies in the rate of modification, which depends upon antigenic affinity. Antibodies with higher affinity are hypermutated at low rate, while, antibodies with a lower affinity are hypermutated at high rate. This phenomenon is called receptor editing and governs the hypermutation that guides the solutions toward local optimum, while, receptor editing helps to avoid local optima.

The main immune aspects taken into account to develop the clonal algorithm are: cloning of the most stimulated cells proportionally to their antigenic affinity; death of non-stimulated cells; affinity maturation and selection of cells proportionally to their antigenic affinity; hypermutation

and generation and maintenance of diversity. Figure 21 clearly shows the flow of clonal algorithm.

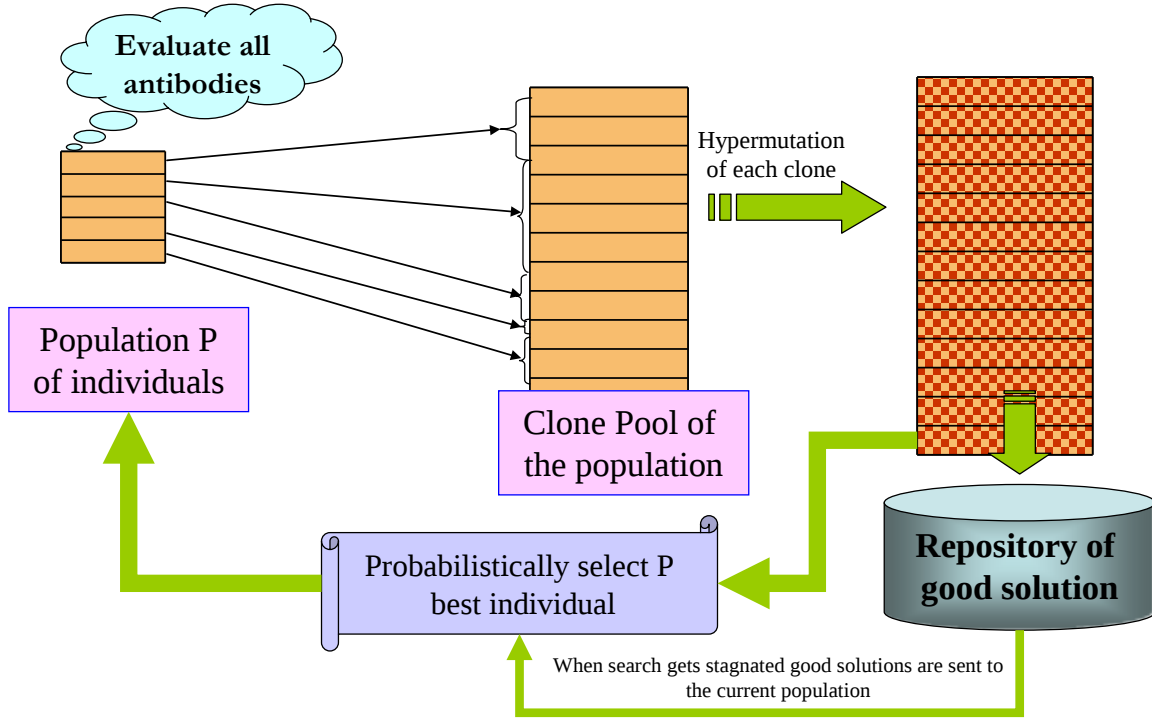


Figure 21: Flow Chart of Clonal Algorithm (Shukla, 2009)

11.5.3 Clonal Algorithm Guided FCM Clustering

It is worth mentioning that FCM works on a calculus based optimization method which can be trapped in local optima during optimization of individual clusters; moreover it is also very sensitive towards the initialization. In order to avoid the entrapment in local optima, and to reduce the sensitivity to initialization, a clonally guided FCM approach has been employed. In this methodology, set of clusters center acts as mandatory string and can be represented as real values or binary values both.

The initial population of size P is constructed by random assignment of real numbers R to each of the S feature of the cluster's center. These values are constrained to be in the range of the feature to which they are assigned. The objective is to optimize the function presented in equation (18) by the aid of clonal algorithm. It has already been mentioned earlier that antibodies (strings) used in clonal algorithm will comprise only prototype V . Therefore, objective function represented in (18) should be reformulated by putting values of μ_{ik} from equation (21) to it. By performing the above procedure, the new objective function takes the following form.

$$R(V) = \sum_{k=1}^c \left(\sum_{i=1}^N \|x_i - v_k\|^{1/(1-m)} \right)^{1-m} \quad (23)$$

Now the approach will be to optimize $R(V)$. Following are the steps involved in propagation of clonally optimized FCM clustering (Shukla et al., 2009):

- Step 0:** Set parameters related to fuzzy clustering- (X: data universe, n: number of data samples, c: desired number of clusters, p: number of feature dimension, m: degree of fuzziness). Set Parameters related to Clonal Algorithm- (max_gen: maximum number of generations, pop_size: population size, Nc: number of clones, HPr: hypermutation rate).
- Step 1:** Initialization- A randomly generated initial population of antibodies depending upon the problem-environment is also required. Generate clusters' center v_{kj} randomly according to range of X.
- Step 2:** Evaluation- Compute the objective function $(R(V) = \sum_{k=1}^c \left(\sum_{i=1}^N \|x_i - v_k\|^{1/(1-m)} \right)^{1-m})$ for entire population.
- Step 3:** Selection- Sort all the individuals of population according to their fitness values and select best antibodies. Fitness value refers to computed objective function $R(V)$ for all individuals.
- Step 4:** Cloning- select the individuals having large antigenic affinity and make higher number of clones for them and less for individuals having low antigenic affinity.
- Step 5:** Hypermutation- Generate a random real number $r \in [0,1]$; if $(r \leq HPr)$ then generate new elements in the j th column of j th individual. Probabilistically select best individuals equal to pop_size.
- Step 6:** Termination- if (number of iteration < max_gen) then Let gen=gen+1, and go to step 2; else terminate the algorithm.

11.5.4 Development of C²FDT

The training data set is clustered in to pre-defined number of clusters so that similar data points are put together. Each cluster is treated as a node of the tree. These entities are subsequently refined according to a selected heterogeneity criterion, named inconsistency index until some termination criterion has met. These subsequent refinements of the clusters result in a tree structure. At any level of the tree, all remaining nodes are allowed to submit their candidature, and the most appropriate among them, i.e., having highest inconsistency index, is chosen for further refinement. This mechanism is schematically shown in Figure 22.

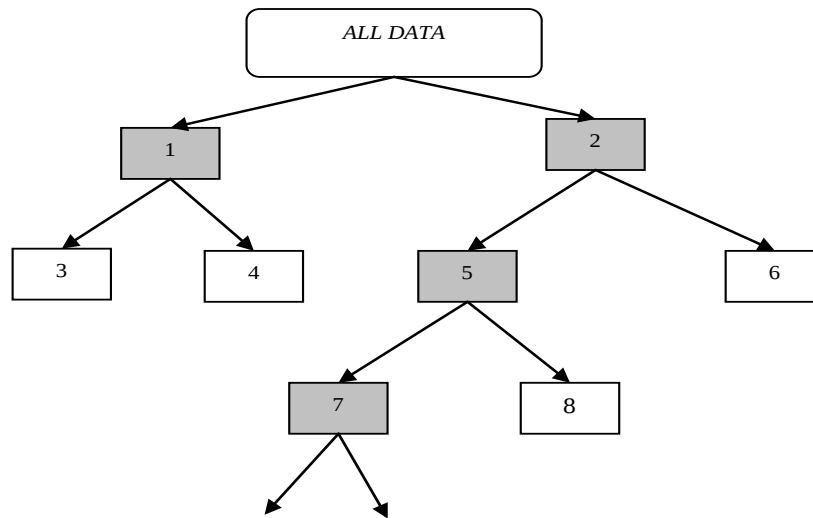


Figure 22: Growth of Decision Tree by Expanding Nodes (Shukla and Tiwari, 2009)

In Figure 22, first and second nodes are clustered from the initial data set. Assuming the first node has larger magnitude of inconsistency index, it is then considered as the appropriate candidate for further splitting. This splitting results two more nodes viz. 3 and 4. Thereafter, nodes 2, 3, and 4 have been compared to obtain the candidate having highest variability. Among these three nodes, node 2 has the highest inconsistency index, thus it is selected for further refinement. In the similar fashion remaining nodes are tested and refined.

Objective of the inconsistency criterion is to quantify dispersion of data allocated to a specific cluster so that higher dispersion of data results in higher values of the criterion. Individual patterns (data points) are associated with a cluster with the different membership grades. However, there exist a cluster in which pattern has maximum degree of membership. Let us represent the i^{th} node of the tree N_i as an ordered triple:

$$N_i = \langle X_i, Y_i, U_i \rangle \quad (24)$$

where X_i refers to all elements of the data set that belong to this node in asset of the highest membership value.

$$X_i = \{x(k) \mid \mu_i(x(k)) > \mu_j(x(k)) \text{ for all } j \neq i\}; \quad (25)$$

Y_i is a set of output coordinates of the elements that have already been assigned to X_i and can be expressed in following manner:

$$Y_i = \{y(k) \mid x(k) \in X_i\}. \quad (26)$$

U_i is membership grade vector of elements in X_i , and is given as:

$$U_i = [\mu_i(x(1)), \mu_i(x(2)), \dots, \mu_i(x(N))] \quad (27)$$

The representative of this cluster positioned in the output space is defined as the weighted sum as follows:

$$\eta_i = \frac{\sum_{(x(k), y(k)) \in X \times Y_{ii}} \mu_i(x(k)) \times y(k)}{\sum_{(x(k), y(k)) \in X \times Y_{ii}} \mu_i(x(k))} \quad (28)$$

The inconsistency of the data in the output space is represented as inconsistency index (Ω), taken as the spread around the representative (η_i). We again consider partial involvement of elements in X_i by weighting distance with the associated membership grade.

$$\Omega_i = \sum_{(x(k), y(k)) \in X \times Y_{ii}} \mu_i(x(k)) (y(k) - \eta_i)^2 \quad (29)$$

For further splitting, we select the node having the highest value of η , such as $\eta_{j_{max}}$, and expend it by forming its c children with the aid of clustering algorithm proposed in this thesis. This process is repeated and nodes of the tree are regularly examined to expand the one with the highest inconsistency index.

Growth of the tree is governed by the clauses under which clusters can further be expanded (Pedrycz and Sosnowski, 2005). This thesis mainly focuses on following two prevailing conditions that tackle the nature of data behind each node. The first one is that, a given node can be split if it consists of ample data points. This number of data must be greater than number of clusters to be formed. Generally, this number is a multiple of c , such as $2c, 3c$ etc (Shukla and Tiwari, 2009). The second stopping criterion is concerned with the structure of the pattern that we are willing to attain through clustering. It is evident that the smaller data results in a less prominent structure. This is reflected by the entries of partition matrix that tend to be equal to each other and equal to $1/c$. If no structure is traced, this equal distribution of membership grades occurs across each column of the partition matrix. This lack of visible structure can be quantified by the following expression:

$$\chi_k = 1 - c^c \prod_{i=1}^c \mu_{ik} \quad (30)$$

If all members of the partition matrix are equal to $1/c$, then the result is equal to zero. If we encounter a full membership to a certain cluster, then the resulting value is equal to 1. For describing the structural dependencies within the entire data set in a certain node, calculations are carried out over all patterns located at the node of the tree.

$$\bar{\chi} = \frac{1}{n} \sum_{k=1}^n \chi_k = \frac{1}{n} \sum_{k=1}^n \left(1 - c^c \prod_{i=1}^c \mu_{ik} \right). \quad (31)$$

Again, with no structurability present in the data, the above expression returns a zero value.

To attain a better feel as to the lack of structure and the ensuring values of equation (31), let us consider a case where, all entries in a certain column of the partition matrix are equal to $1/c$ with the slight deviation equal to ε . In 50 percent of the cases, it is considered that these entries are higher than $1/c$, and we put $\mu = 1/c + \varepsilon$; in the remaining 50percent, we consider the decreased over $1/c$ and have $\mu = 1/c - \varepsilon$. Furthermore, let us treat ε as a fraction of the original membership value, that is, makes it equal to $\lambda(1/c)$, where $\lambda \in [0, 1/2]$. Then, equation (31) becomes:

$$\chi = 1 - c^c (1/c + \varepsilon)^{c/2} (1/c - \varepsilon)^{c/2} = 1 - c^c \left(\left(\frac{1}{c} (1 + \lambda) \right) \right)^{c/2} \left(\frac{1}{c} (1 - \lambda) \right)^{c/2} \quad (32)$$

Two above mentioned measures can be used as a stopping criterion in the growth of the tree. A node may be assumed obstinate once the number of pattern falls under the assumed threshold and/or the structurability index is too low. The first index comes in the form of precondition; if not satisfied, it prevents us from expanding the node. The second index is a sort of some post condition; to evaluate its value, we have to cluster the data first then determine its value. These two indexes guide the decision making that focuses on where to split the data. It does not guarantee that the resulting tree will be the best from the point of view of classification or prediction of continuous output variable. The criterion of diversity (sum of Ω_i at the leaves) can also be used as termination criterion. Another possible termination option is to monitor their changes along the node of the tree once build.

11.5.5 Use of C²FDT in Forecasting Problems

C²FDT can be used to forecast dependent variable ($\hat{y}$) associated with input vector (x). In the calculation of ($\hat{y}$), we make use the membership grades computed for each clusters as follows:

$$\mu_i(k) = \frac{1}{\sum_{j=1}^c \left(\frac{\|x - v_i\|}{\|x - v_j\|} \right)^{2/(m-1)}} \quad (33)$$

where $\|x - v_i\|$ is a distance measure between x and v_i . These calculations are carried out to the leaves of the tree, by calculating $\mu_i(x)$ for each node of the tree and then selecting the corresponding path and moving down. This process is shown in Figure 23.

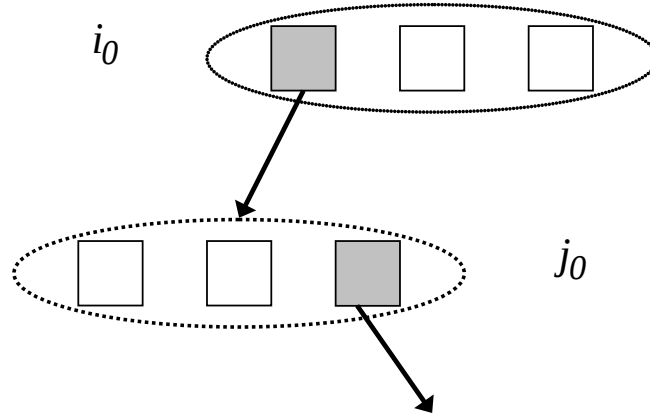


Figure 23: Traversing a Decision Tree

At some level, path i_0 is determined where $i_0 = \arg \max_j \mu_j(x)$. This process is repeated for each level of the tree and predicted value of occurring at the final leaf node is equal to η_i as given in equation (28).

It becomes important to show the boundaries of the classification regions produced by the clusters and contrast them with the geometry of classification regions generated by decision trees. The following criterion is used for assignment of patterns to the clusters. Assign x to class ω_i if, at the same level, $\mu_i(x)$ exceeds the value of the membership in all remaining clusters, that is, $\mu_i(x) > \max_j \mu_j(x)$.

12.0 SOFTWARE DEVELOPMENT

The methodology discussed in this research includes various artificial intelligence techniques integrated with a database management system to predict the future workforce for an organization. Therefore, for a successful implementation of the methodology, an in-depth knowledge of Artificial Intelligence, Human Resource Management, Database Management Systems, Software development and Graphical User Interface development is required. Moreover, there arises a need of integrated software which has a simple yet powerful frontend for operating all the features of the forecasting software.

As one of the deliverables of the project, workforce forecasting software has been developed as a part of this research. The software utilizes various high level programming languages and the object oriented programming techniques for its successful implementation. It also makes use of the standard Microsoft Windows Graphical User Interface to make the job easier for the end user while operating the software. Following tools have been used to design the workforce forecasting software: *MATLAB R2010a*, *MS Excel 2007*, and *Visual C#*.

MATLAB R2010a (Figure 24) has been utilized as the programming language to write the functions which can be later utilized for making the dynamically linked libraries (.DLL files). The three functions that have been developed within the MATLAB are as follows:

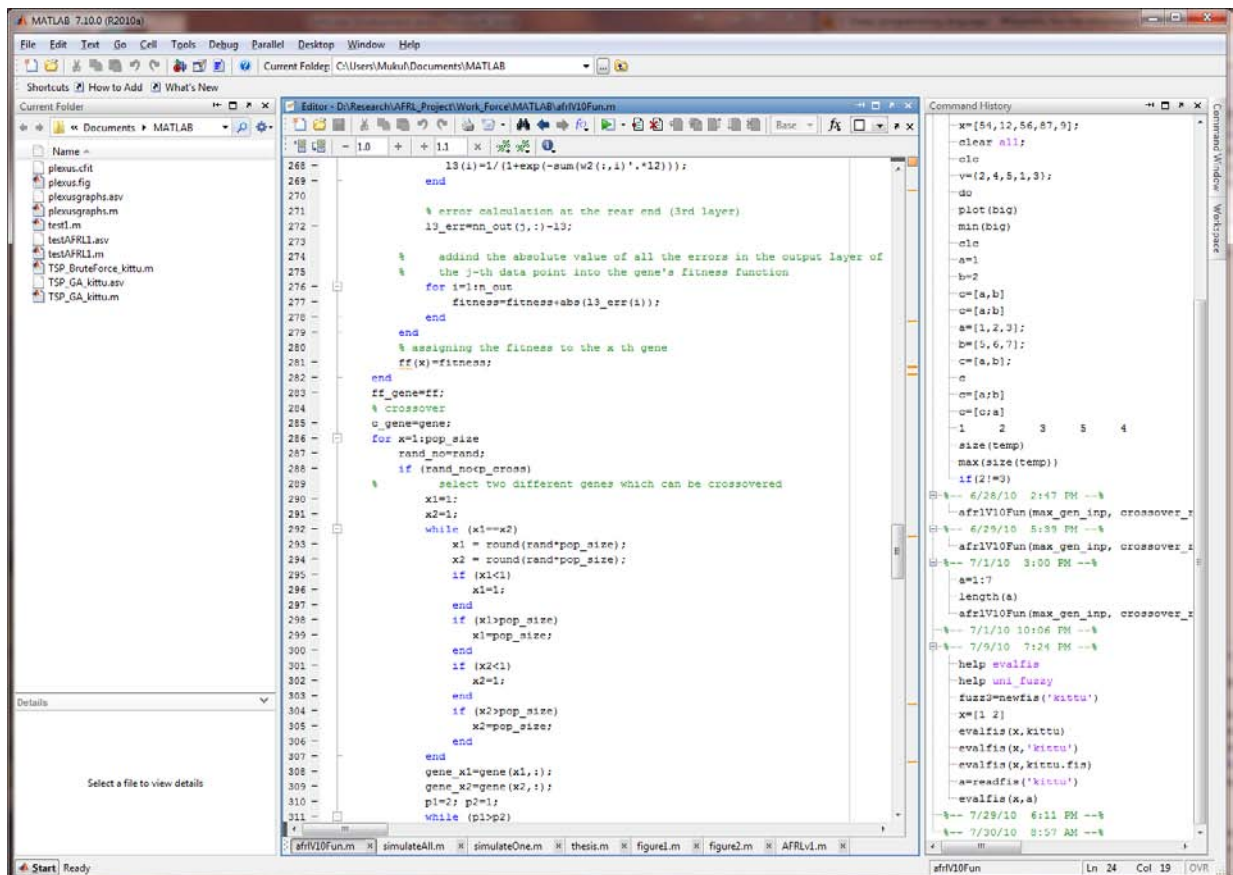


Figure 24: MATLAB Standard IDE for Software Development

- 1) **Function “AFRLv10Fun.m”** - This is the main function which takes in 8 inputs and produces 8 outputs as shown in Figure 25.

```
function [error_in, error_final, ret_gen, ret_n_inp, ret_n_hid, ret_n_out,
        ret_w1, ret_w2] = afrlv10Fun(max_gen_inp, crossover_rate,
        mutation_rate, population_size, nn_hidden, seed_no, err_supply, cond)
```

Figure 25: MATLAB Function “AFRLv10Fun.m” Structure

As shown in Figure 25, AFRLv10Fun.m takes 8 inputs as

1. max_gen_inp: Max Generation used as the stopping criteria.
2. crossover_rate: Initial crossover rate provided to the Genetic Algorithm
3. mutation_rate: Initial mutation rate provided to the Genetic Algorithm
4. population_size: Population size provided to the Genetic Algorithm
5. nn_hidden: Number of neurons in the hidden layer
6. seed_no: Initial Random Seed
7. err_supply: Stopping criteria for minimum error supply
8. cond: Tells us whether stopping criteria is 1 or 7 or both

Similarly, the outputs in the AFRLv10Fun.m are listed as:

1. error_in: Input error value required by the user as the stopping criteria
2. error_final: Final Error value attained by the program
3. ret_gen: Number of generations before the final error value was generated.
4. ret_n_inp: Number of input in the input nodes.
5. ret_n_hid: Number of neurons in the hidden layers.
6. ret_n_out: Number of neurons in the output nodes.
7. ret_w1: Weights within the neural network
8. ret_w2: Weights within the neural network

Besides, this function also utilizes a set of commands which facilitate the selection of the database file. These set of commands have been shown below in Figure 26.

```
counter=0;
while(counter==0)
    [FileName,PathName] = uigetfile('data*.xls','Select the Data File');
    if(FileName==0)
        counter=0;
    else
        counter=1;
    end
end
file=fullfile(PathName,FileName);
```

Figure 26: MATLAB Code for Selecting the Data File

- 2) **Function “SimulateAll.m”** – Figure 27 shows the implementation of another function within the AFRL software which generates the results by simulating the optimized neural network and thereafter, compares it with the original results within the dataset.

```

MATLAB 7.10.0 (R2010a)
File Edit Text Go Cell Tools Debug Parallel Desktop Window Help
Current Folder: D:\Research\AFRL_Project\Work_Force\MATLAB
Current Folder: Work_Force \ MATLAB
Name =
er03999.mat
For_Presentation.asv
For_Presentation.m
For_Presentation2.asv
For_Presentation2.m
inputs
integrateCFDT.asv
integrateCFDT.m
lab.m
mod_check.m
mutation_change_name
NewDataSet.asv
NewDataSet.m
newSmallData.m
pi_graph.asv
pi_graph.m
plot1.m
progressBar.asv
progressBar.m
sanjay1.m
sanjay2.asv
sanjay2.m
saveExcel.m
simulate.asv
simulateAll.asv
simulateAll.m
SimulateAllLib.pj
simulateOne.asv
simulateOne.m
SimulateOneLib.pj
test1_no_inp.m
testCFDT.asv
testCFDT2.asv
testCFDT2.m

simulateAll.m (MATLAB Function)
the company f gives its important factor as following
simulateAll(inp,n_hid,n_out,w1,w2)

95 p(14)=q(27);
96
97 p(15)=(q(28)+ q(29)+ q(30)+ q(31)+ q(32)+ q(33)+ q(34)+ q(35)+ q(36)+ q(37)+
98
99 p(16)=(q(44)+q(45))/2;
100
101 y(1,1)=round(p(1)+0.8*(180+round(rand*20))+p(2)+0.6*(190+round(rand*20))+ p(3
102
103 y(1,2)=round(p(1)+0.67*(80+round(rand*20))+p(2)+0.345*(90+round(rand*20))+ p(
104
105 p_all=[p_all;p];
106
107 end
108
109 nn_inp=p_all;
110 max_wf=2000;
111 nn_out=y/max_wf;
112
113 final=[];
114 for j=1:n_data
115     l1=nn_inp(j,:);
116
117     for i=1:n_hid
118         l2(i)= 1/(1+exp(-sum(w1(:,i),'*l1')));
119     end
120
121     for i=1:n_out
122         l3(i)=1/(1+exp(-sum(w2(:,i),'*l2')));
123     end
124
125     final=[final;l3*max_wf];
126 end
127 % To check the plot of the simulated data
128
129 a=plot(1:n_data, final(:,1), 1:n_data, nn_out(:,1)*max_wf);
130 title('Workforce Category #1: Green is the Input and Blue is the Simulated Output');
131 xlabel('Data');
132 ylabel('Number of Employees');
133 b=figure;
134 b=plot(1:n_data, final(:,2), 1:n_data, nn_out(:,2)*max_wf);
135 title('Workforce Category #2: Green is the Input and Blue is the Simulated Output');
136 xlabel('Data');
137 ylabel('Number of Employees');
138 end

Command History
x=[54,12,56,87,9];
clear all;
c1o
v=[2,4,5,1,3];
do
plot(b1q)
min(b1q)
c1o
a=1
b=2
c=[a,b]
a=[1,2,3];
b=[5,6,7];
c=[a,b];
c
a=[a;b]
c=[c;a]
1 2 3 5 4
size(temp)
max(size(temp))
if(2==3)
6/28/10 2:47 PM --%
afzV10Fun(max_gen_inp, crossover_r
6/29/10 5:33 PM --%
afzV10Fun(max_gen_inp, crossover_r
7/1/10 3:00 PM --%
a=1:7
length(a)
afzV10Fun(max_gen_inp, crossover_r
7/1/10 10:04 PM --%
7/9/10 7:24 PM --%
help evalfis
help uni_fuzzy
fuzzy=newfis('kittu')
x=[1 2]
evalfis(x,kittu)
evalfis(x,kittu.fis)
evalfis(x,kittu.fis)
a=readfis('kittu')
evalfis(a,a)
7/29/10 6:11 PM --%
7/30/10 8:57 AM --%

```

Figure 27: Function “SimulateAll.m” Implementation

- 3) **Function “SimulateOne.m”** - The third function of interest is the SimulateOne.m which is shown in Figure 28. It takes up one value for the future input and generates the result for the future.

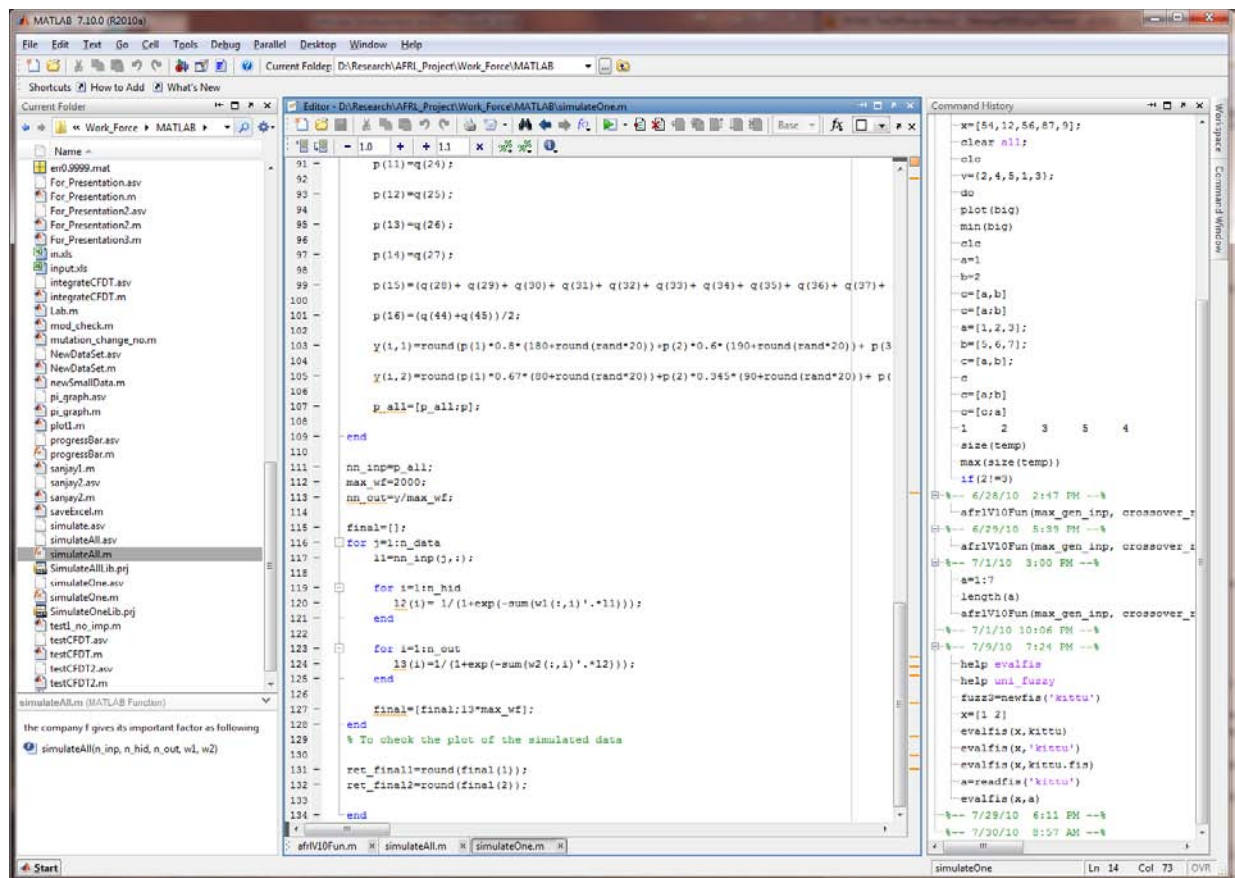


Figure 28: Function “SimulateOne.m” Implementation

After these functions have been coded in MATLAB, they can be utilized by the programmer to predict the future workforce. However, these functions do not provide any graphical user frontend. Keeping this in mind, a graphical user interface has also been developed to facilitate easy access to these functions and to maintain the database at the same time. For this purpose, Visual C# is utilized to make the GUI. However, before using these functions into a Visual C# these MATLAB functions are compiled using the MATLAB compiler as shown in Figure 29. Deployment tool compiles the MATLAB functions into dynamically linked libraries (DLL) which can be further used into the Visual C#.

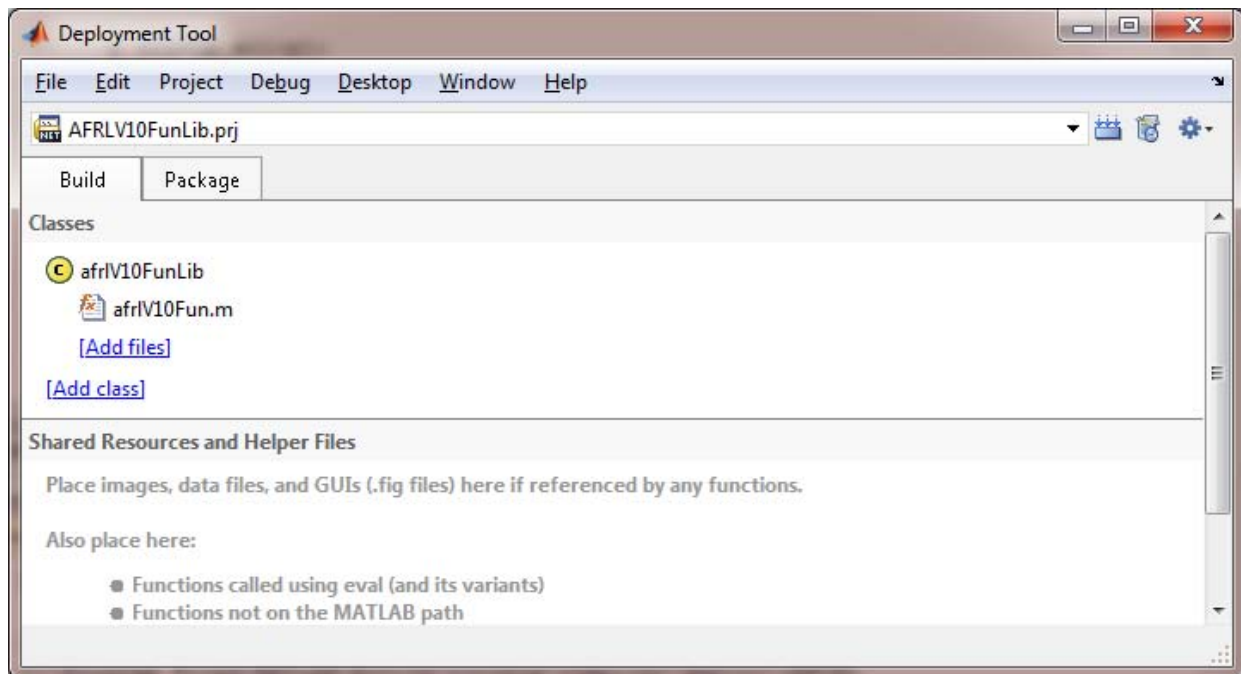


Figure 29: Deployment Tool

By using the deployment tool in MATLAB, three DLL files are generated corresponding to each MATLAB function.

Once the DLL files are obtained, the final task is to incorporate them into the Visual C# program for developing the graphical user interface. Figure 30 shows the solution explorer for the software where we can see several references in the list. A few of these libraries need more description as below:

- **AFRLV10FunLib, SimulateAllLib, SimulateOneLib:** These are the libraries containing the AFRLV10Fun, SimulateAll and SimulateOne as discussed before.
- **Excel:** This is the library for accessing the database stored in the Excel.
- **MWArray:** This library that lets Visual C# interact with the MATLAB array types.

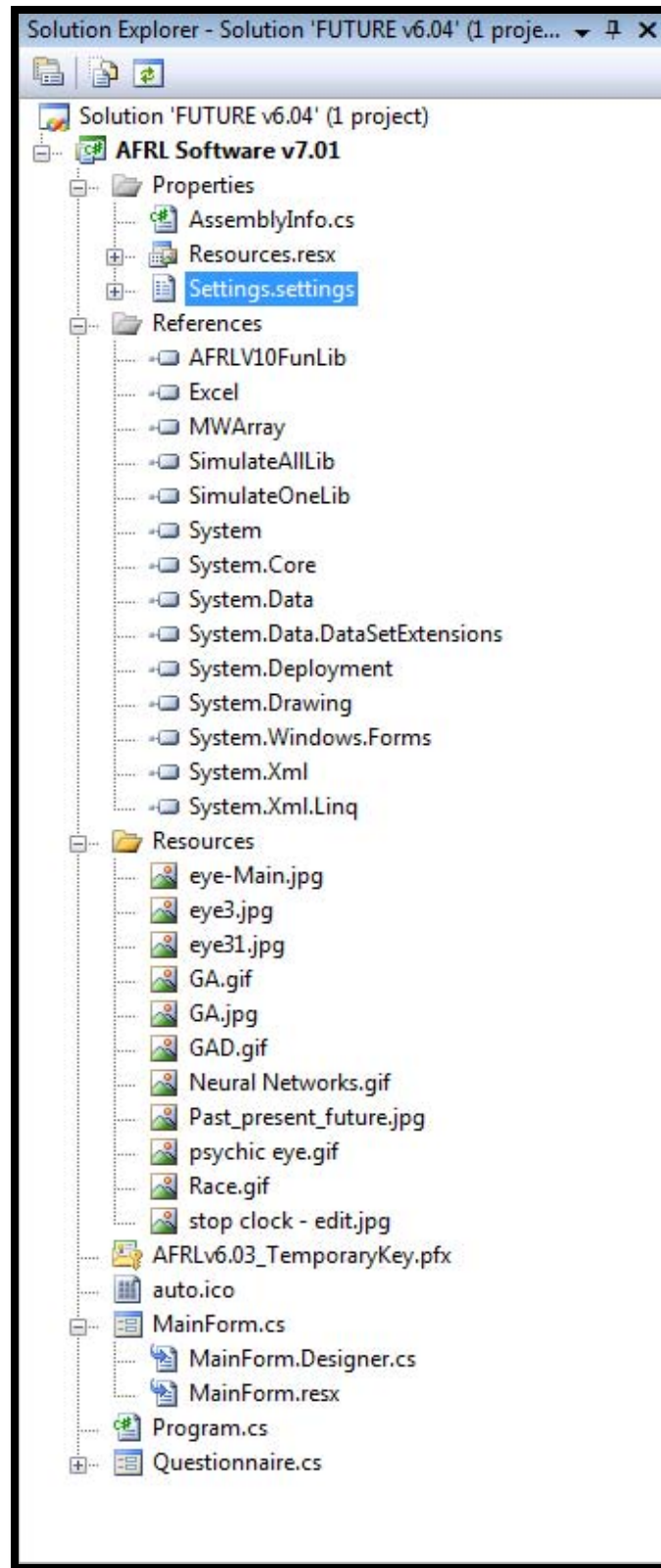


Figure 30: Solution Explorer for the Software

Thereafter, the forms are designed in the Visual C# to present the GUI software. Each form in Visual C# has an embedded code running in the background. Figure 31 shows the Mainform.cs used to build up the basic interface for the software.

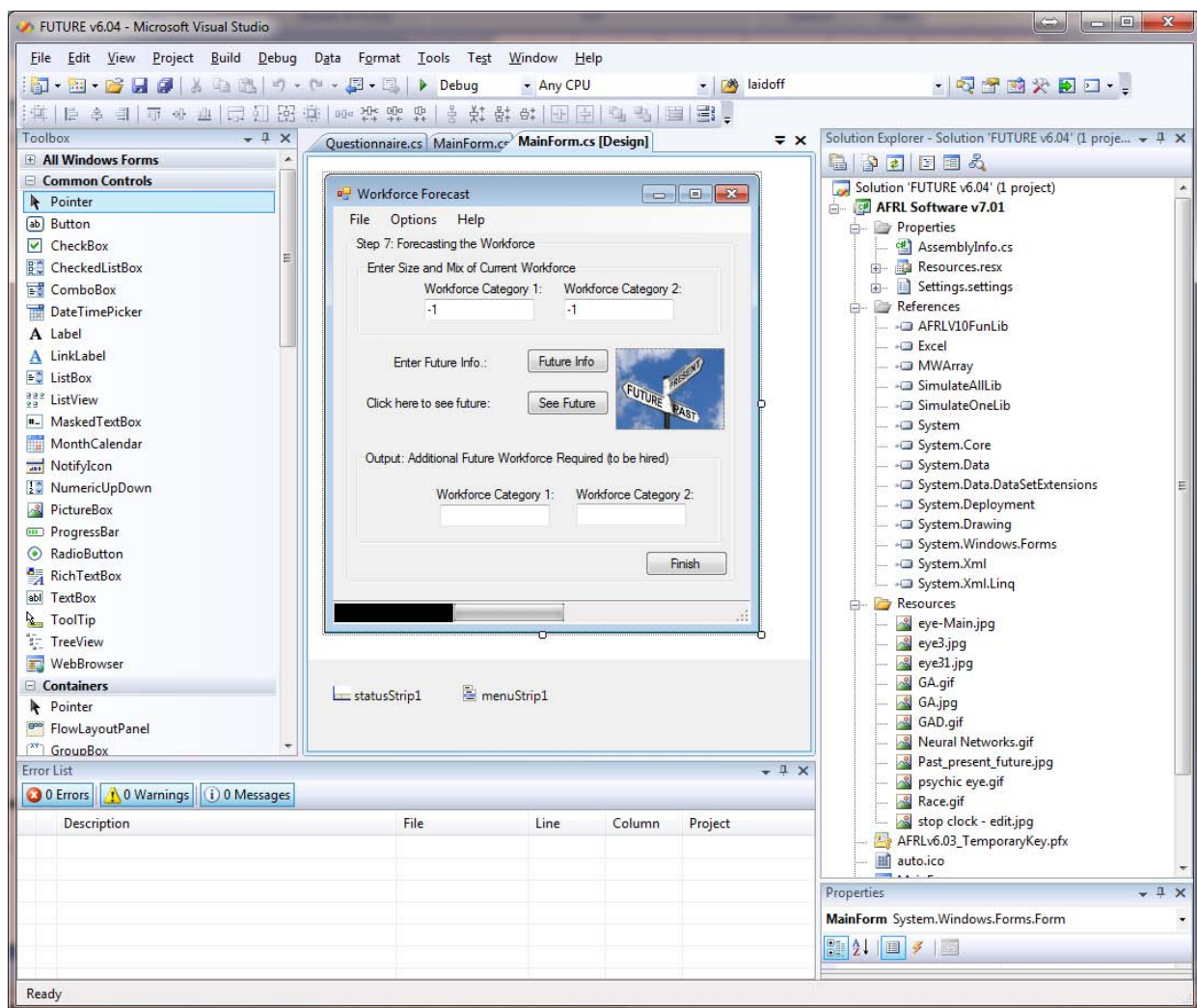


Figure 31: “Mainform.cs” Design Implementation

The coding part of the Mainform.cs is shown in Figure 32.

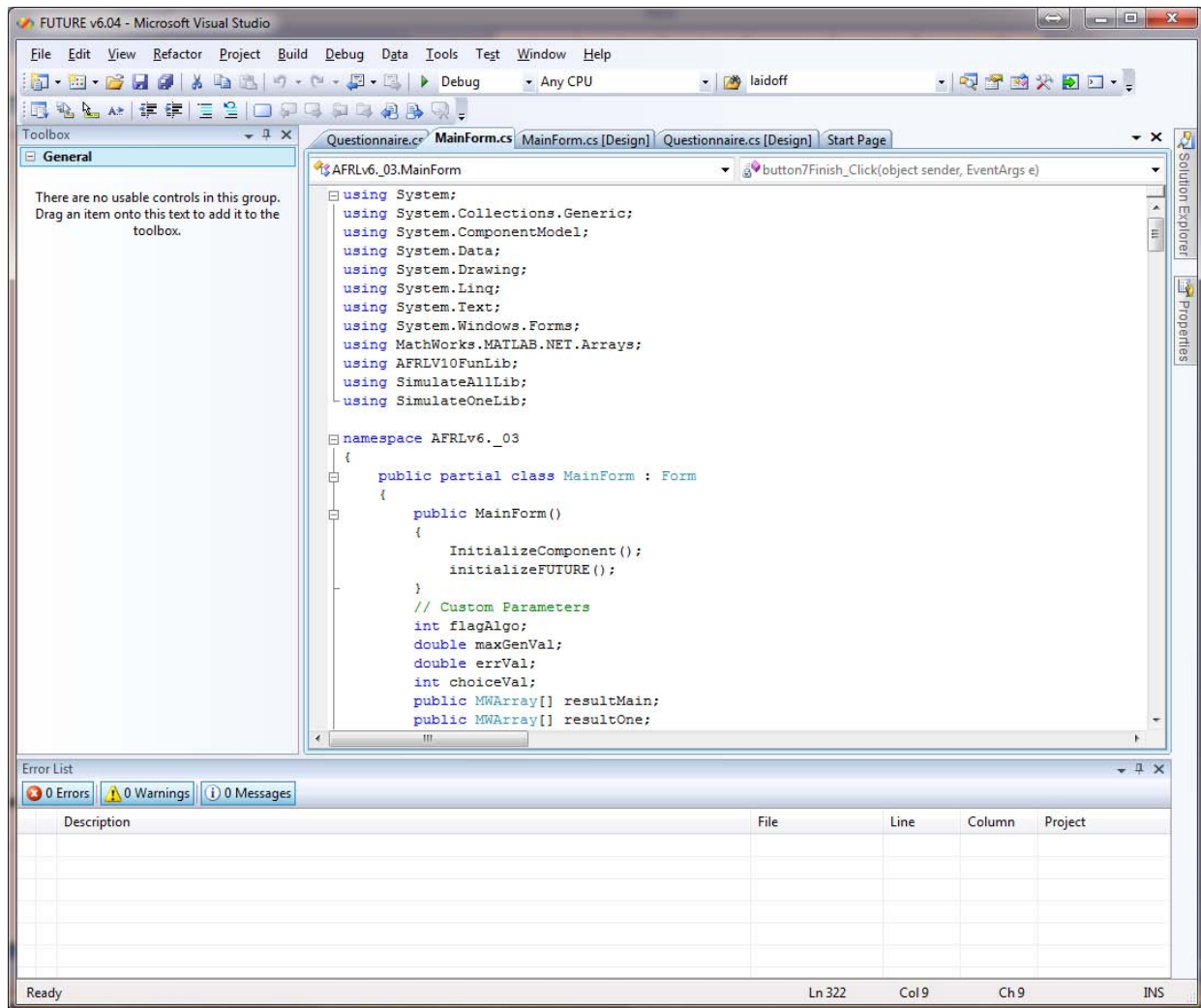


Figure 32: Coding of “Mainform.cs”

Besides, there is secondary form which helps in the easy database management. This is shown in Figure 33.

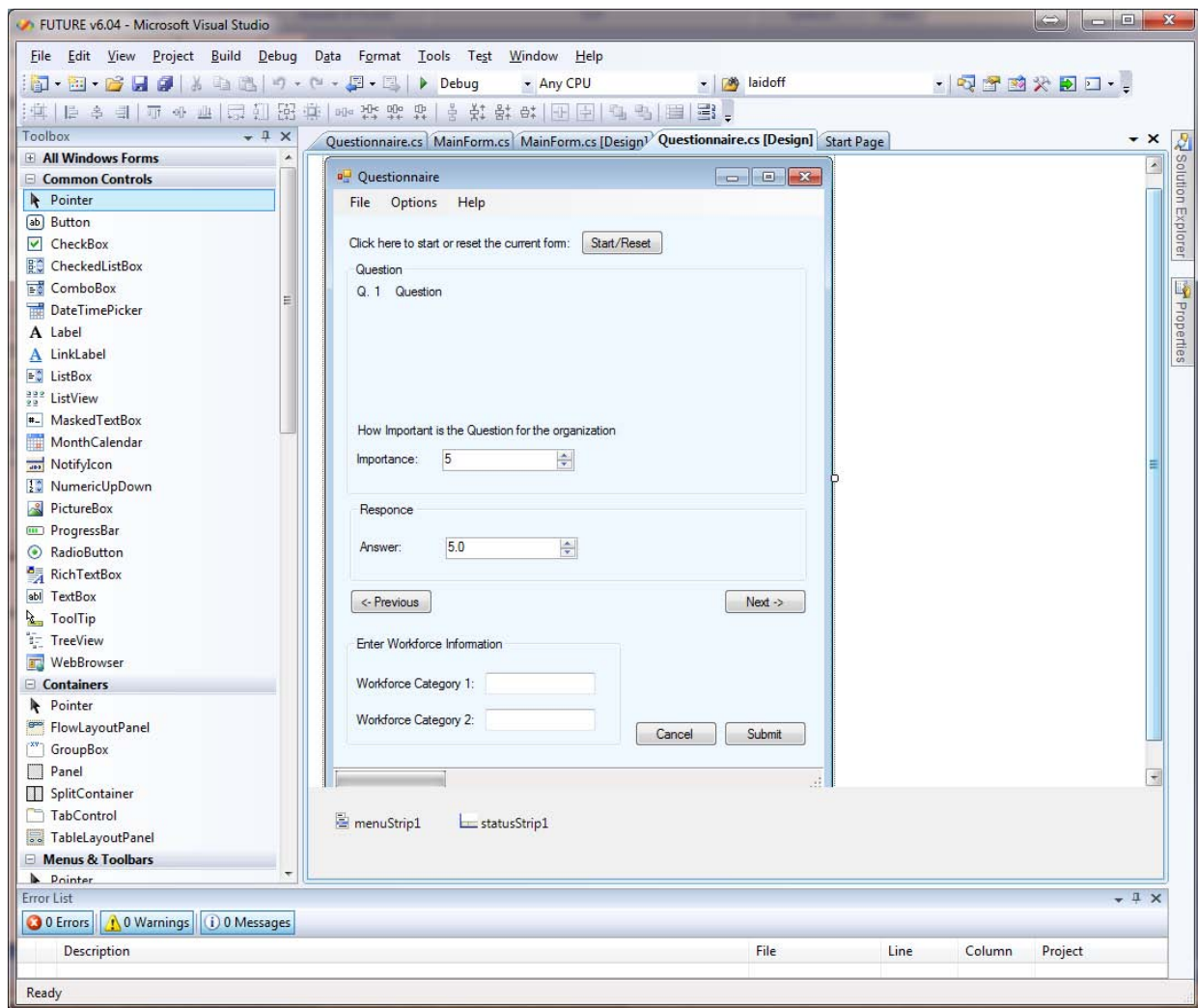


Figure 33: “Questionnaire.cs” Form Implementation

The coding behind the Questionnaire.cs is shown in Figure 34.

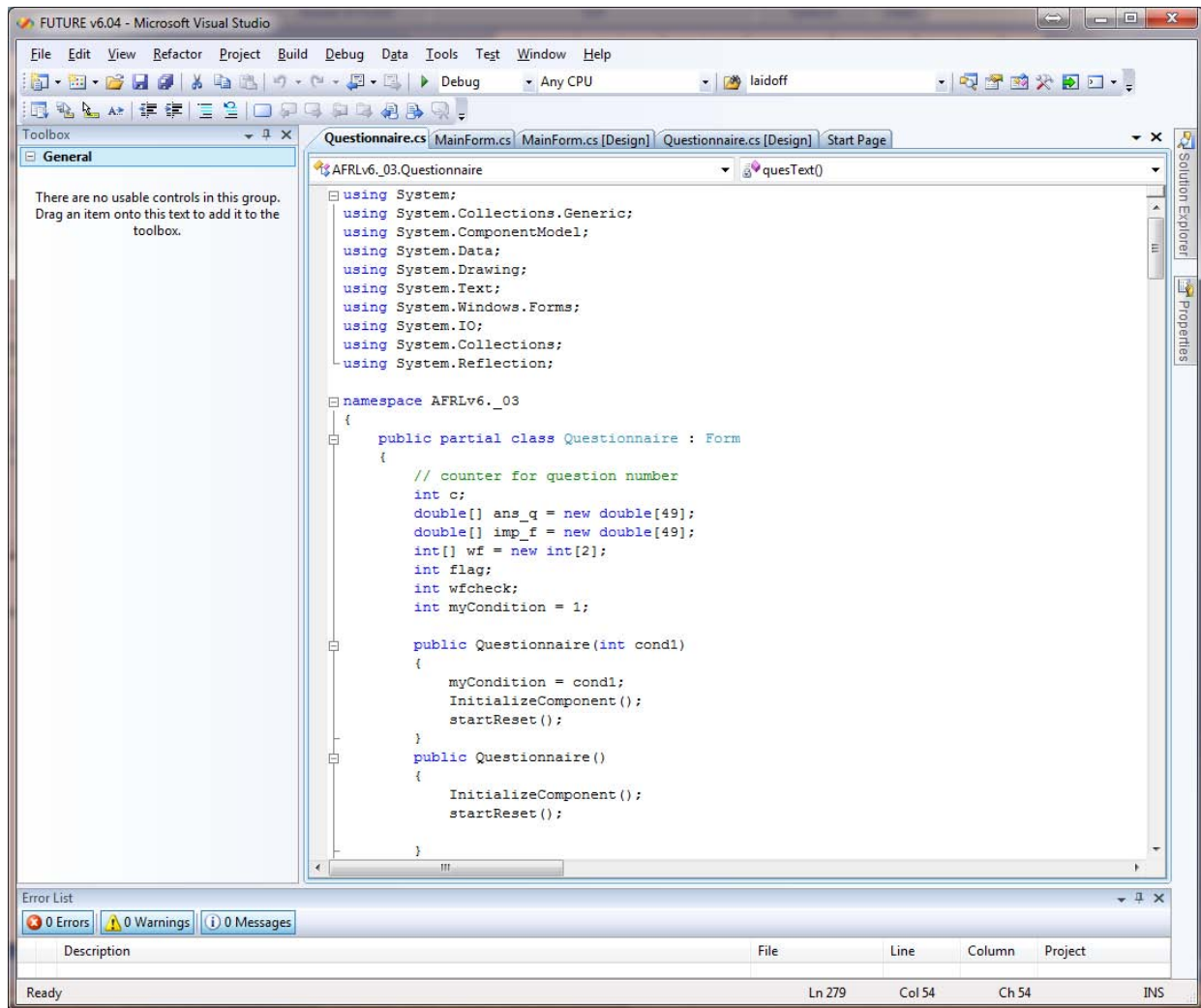


Figure 34: Coding of the “Questionnaire.cs”

Once the forms and the appropriate coding have been done, the final step is to publish the software as shown in Figure 35. Thereafter the final software is obtained in the executable format which is ready to be distributed.

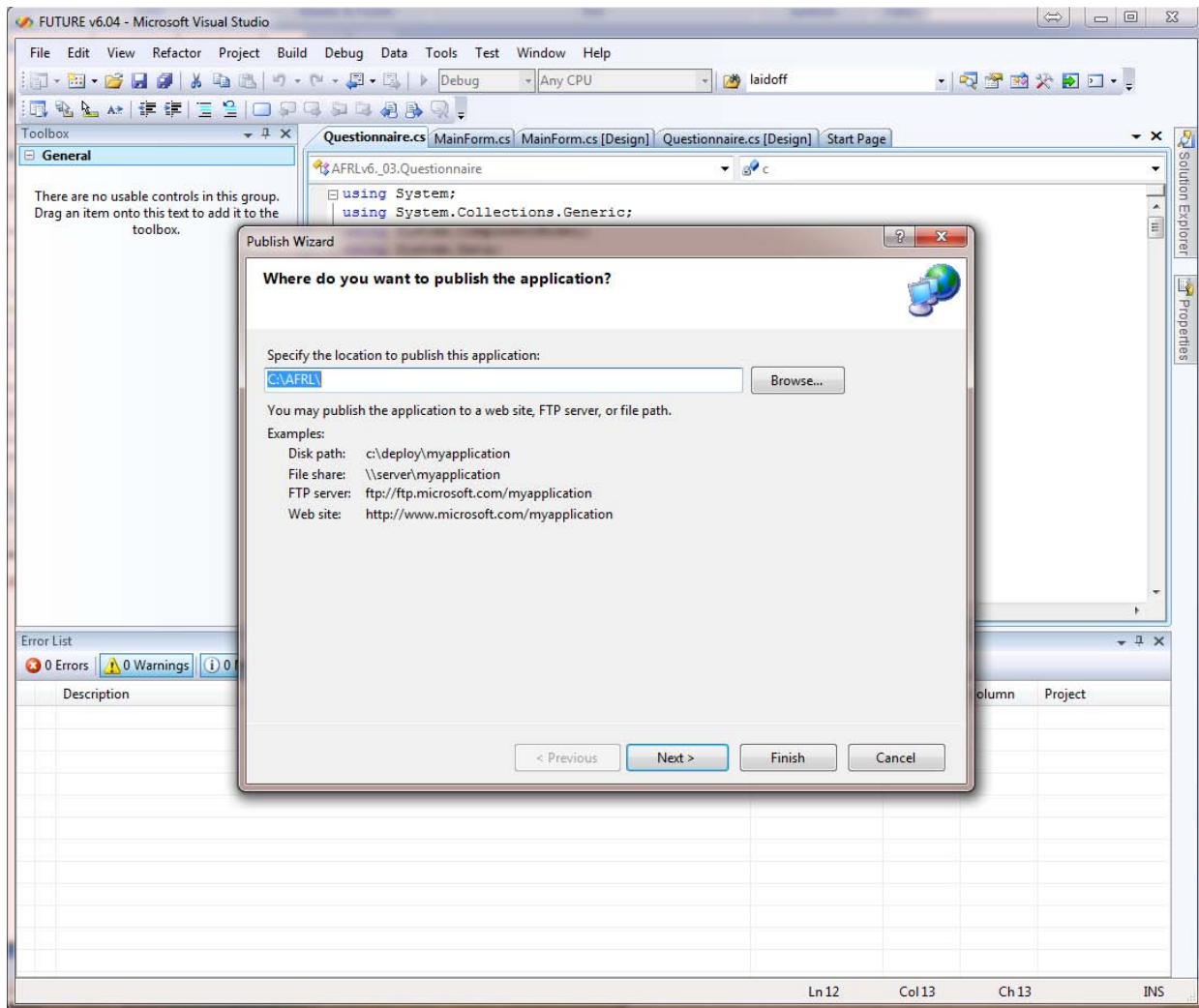


Figure 35: Publish Application

Please refer to the *Quick Start Guide* to learn how to operate the software.

13.0 RESULTS AND DISCUSSION

13.1 Overview of the Simulated Case Study

A simulated case study of a hypothetical organization has been undertaken to test the efficacy of the proposed algorithm and its supremacy over the C²FDT. This hypothetical organization needs to predict its workforce information for the coming 5 years. The first step challenge for the firm to achieve this goal is to set up an expert team for the workforce forecasting. This workforce forecasting team should contain experts from various departments and is also responsible for setting up organization's strategic goals. The ultimate objective of the team is to forecast size and mix of workforce for next coming years to accomplish these goals. In this research, mix of the workforce constitutes more than one category to demonstrate that the proposed methodology is capable of working with several categories of employees within an organization (in this case category 1 and category 2).

The software has been coded and compiled in MATLAB and Microsoft Excel has been utilized for managing data. The graphical user interface (GUI) for the software has been made in Visual C#. For machines which do have MATLAB installed require latest version of MATLAB component runtime (MCR) installed to run the software.

13.2 Survey Questionnaire

After setting the organization's strategic goals, the next task is to conduct internal and external environmental scanning. This scanning helps in identifying the parameters that govern the size and mix of workforce. In this research, we have simulated the internal and external environmental scanning by conducting computer simulation for the survey questionnaire. Figure 36 shows the GUI implementation of the actual survey form.

Figure 36: The Survey Form

As stated before, the answer to each question is twofold: (1) the actual answer (2) the importance value of the question. Figure 36 shows the graphical frontend for collecting the information. This information is saved in the Excel file for later use.

Thereafter, these questions are grouped into 17 parameters which were shown in Table 5. The relationship between parameters and the questions (M_{ij}) is shown in a separate Table 6 as shown below. Mathematically,

$$M_{ij} = \begin{cases} 1 & \text{if question contributes to } x_j \\ 0 & \text{otherwise} \end{cases} \quad (1)$$

Table 6: Relationship Matrix between Parameters and Questions

	Parameters																		No.
	No.	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	
Q u e s t i o n s	1	1	0	0	0	0	0	1	1	1	1	0	0	0	0	0	0	0	1
	2	1	0	0	0	0	0	0				0	0	0	0	0	0	0	2
	3	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	3
	4	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	4
	5	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	5
	6	0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	6
	7	0	0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	7
	8	0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	8
	9	0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	9
	10	0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	10
	11	0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	11
	12	0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	12
	13	0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	13
	14	0	0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	14
	15	0	0	0	0	1	0	0	0	0	0	0	0	0	0	0	1	0	15
	16	0	0	0	0	1	0	0	0	0	0	0	0	0	0	0	1	0	16
	17	0	0	0	0	1	0	0	0	0	0	0	0	0	0	0	1	0	17
	18	0	0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	18
	19	0	0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	19
	20	0	0	0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	20
	21	0	0	0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	21
	22	0	0	0	0	0	0	1	0	0	0	0	0	0	0	0	0	0	22
	23	0	0	0	0	0	0	0	1	0	0	0	0	0	0	0	0	0	23
	24	0	0	0	0	0	0	0	0	1	0	0	0	0	0	0	0	0	24
	25	0	0	0	0	0	0	0	0	0	1	0	0	0	0	0	0	0	25
	26	0	0	0	0	0	0	0	0	0	0	1	0	0	0	0	0	0	26
	27	0	0	0	0	0	0	0	0	0	0	0	1	0	0	0	0	0	27
	28	0	0	0	0	0	0	0	0	0	0	0	0	1	0	0	0	0	28
	29	0	0	0	0	0	0	0	0	0	0	0	0		1	0	0	0	29
	30	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	0	30
	31	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	31
	32	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	32
	33	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	33
	34	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	34
	35	0	0	0	0	0	0	1	0	0	0	0	0	0	0	0	1	0	35
	36	0	0	0	0	0	0	1	0	0	0	0	0	0	0	0	1	0	36
	37	0	0	0	0	0	0	1	0	0	0	0	0	0	0	0	1	0	37
	38	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	38
	39	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	39
	40	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	40
	41	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	41
	42	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	42
	43	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	43
	44	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	44
	45	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	45
	46	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	46
	47	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	47
	48	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	48
	49	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	49
	50	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	50
	51	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	51
	52	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	52

13.3 The Fuzzy Logic Controller

After collecting the sufficient data, next task is to amalgamate importance of questions and actual answers. In order to accomplish this task, decision support system switches to third phase. In this phase, both the responses (importance and actual answers to the questions) are amalgamated into one pooled performance measure with the aid of Fuzzy Logic Controller (FLC).

FLC works in following three steps: fuzzification, fuzzy rules firing, and defuzzification. Importance to questions and actual answers work as inputs for FLC. In first step, the two variables are fuzzified as shown in Figure 37. For the fuzzification of two inputs, Gaussian function is used and, outcome of this process is membership values of these two inputs in fuzzy sets (Low, Medium and High). After completing fuzzification, next task is firing of rules on membership values of two fuzzified inputs and their linguistic variables. In order to accomplish this task, we have developed a set of fuzzy rules as shown in Figure 37. In the figure it can be seen that fuzzy rules work in following manner: if input 1 is low and input 2 is low then output is small. Rule firing process will result in amalgamation of two inputs in one fuzzy output with associated linguistic variable. After accomplishing the rule firing process next task is to convert the fuzzy amalgamated output into crisp output (in range 0-1). Figure 37 shows the MATLAB Fuzzy Logic Toolbox customized for the problem at hand.

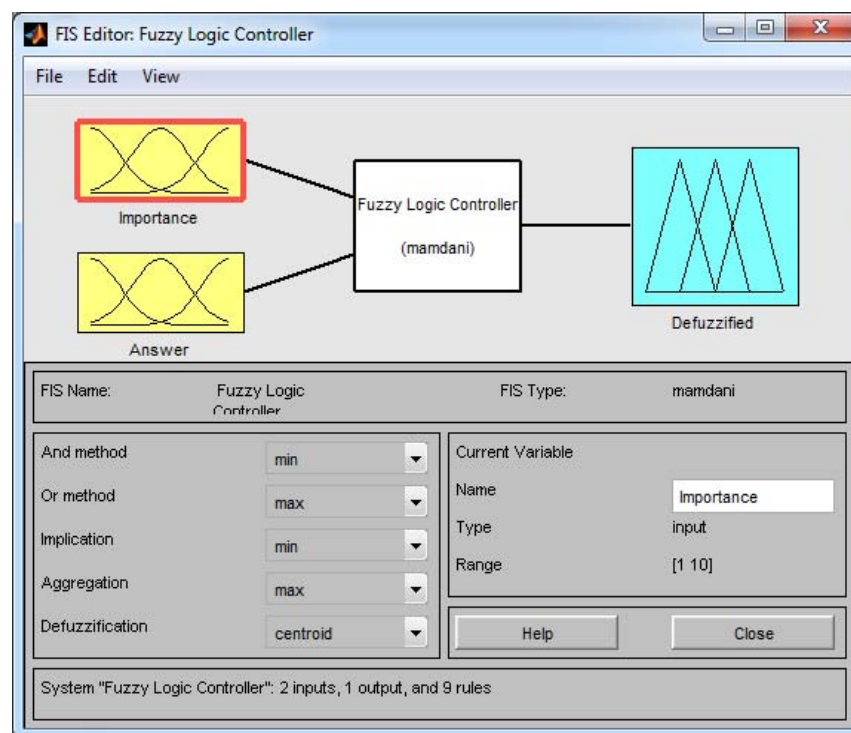


Figure 37: Fuzzy Logic Controller Setup (MATLAB Fuzzy Logic Toolbox)

Figure 38 shows the defuzzified output corresponding to the Importance and the answer values. Similarly, Figure 39 shows the fuzzy rule-base associated with the workforce forecasting problem.

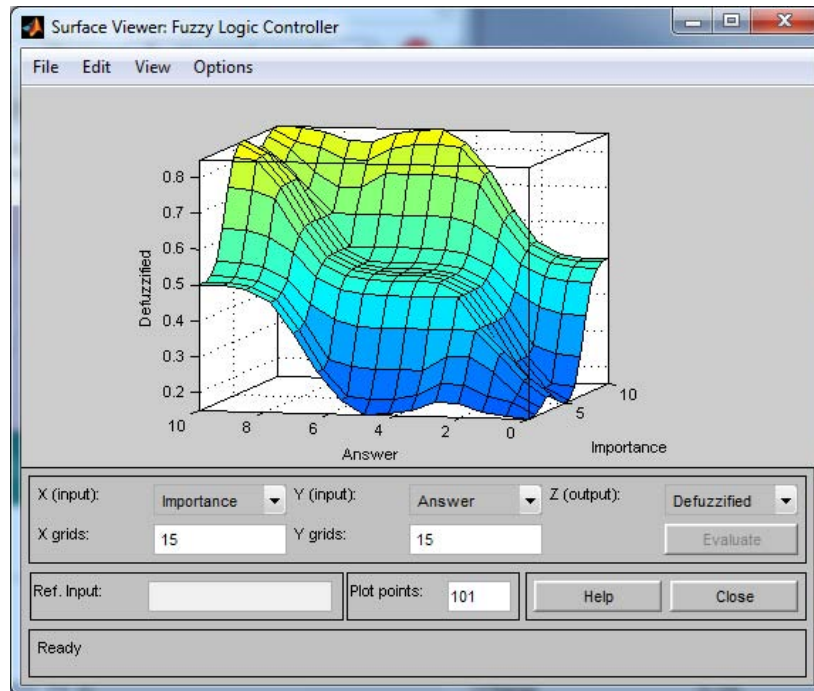


Figure 38: Defuzzified Output Corresponding to the Answer and Importance Value

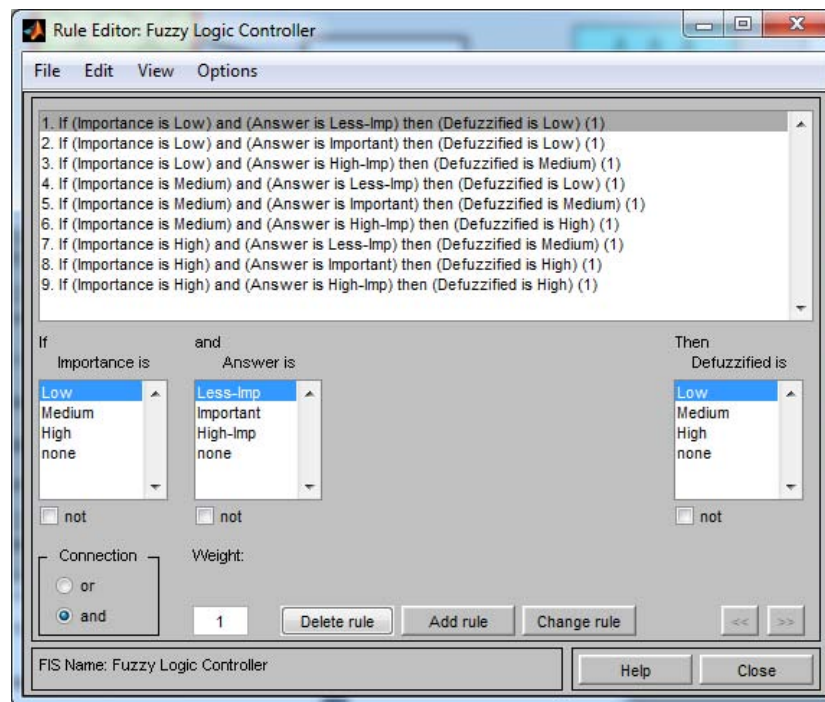


Figure 39: The Fuzzy Rule Base

Now, by utilizing the Fuzzy Logic Controller as shown in Figure 37, we can calculate the defuzzified outputs for each question. Each defuzzified output is referred to as a *partial*

quantified parameter (pqi). Now by utilizing Table 6 we calculate the *actual parameter* (x_j) value as shown in the following equation:

$$x_j = \frac{\sum_{i=1}^{52} pqi \cdot M_{ij}}{\sum_{i=1}^{52} M_{ij}} \quad (35)$$

where, $j=1,2,3\dots 17$ and $i=1,2,3\dots 52$.

Thereafter, two random functions are employed to generate the number of employees in the i^{th} category. By adopting this approach, we are able to formulate a complex function with input variables and certain level of randomness inbuilt in it.

$$y_i = \sum_{j=1}^{17} R_{ij}(x_j) \quad (36)$$

where $R_{ij}(x_j)$ is the random function such that y_i is generated under controlled limits. Thus for two employee categories, we have $i=1, 2$. Table 7 shows the input output dataset format and utilizes above equation to simulate a relationship between the variable and the workforce.

Table 7: The Input Output Dataset Format

Attributes	Data Point 1	Data Point 2	Data Point 3	Data Point 4	Data Point 5
X1	0.5	0.6312879	0.5	0.82564529	0.14461752
X2	0.32652268	0.69212338	0.40474007	0.42414372	0.50029609
X3	0.49293992	0.6102567	0.43674192	0.42121578	0.50839377
X4	0.61991837	0.58241343	0.52858559	0.66223303	0.50006586
X5	0.32881441	0.51035859	0.32874362	0.72421141	0.67081403
X6	0.82194652	0.82014685	0.82194652	0.5	0.5
X7	0.33078948	0.22078771	0.47957535	0.17435471	0.77921229
X8	0.82194652	0.52130935	0.5	0.82564529	0.33078948
X9	0.17805348	0.86293742	0.33078948	0.47869065	0.47957535
X10	0.17805348	0.82564529	0.52130935	0.52130935	0.82194652
X11	0.85538248	0.33078948	0.17805348	0.52130935	0.47957535
X12	0.5	0.33235643	0.17805348	0.66921052	0.47957535
X13	0.5	0.17985315	0.17435471	0.86293742	0.47869065
X14	0.33235643	0.5	0.6312879	0.47869065	0.66921052
X15	0.55068955	0.54795277	0.5083751	0.52105241	0.4184236
X16	0.5	0.5971787	0.40552354	0.576368	0.5
X17	0.60671562	0.55768268	0.42991258	0.54380889	0.49980261
W/F 1	911	1021	781	1167	897
W/F 2	432	496	398	531	457

Moreover, the organizations future information is also recorded by filling out the survey form for the forecasted year which serves as the corresponding input for the forecasted year. In this way, around 125 data points are collected for testing and validation purposes.

13.4 Implementing SGA GONN on the Dataset

After enough data points have been calculated, the algorithm is ready for input data. However, before executing the algorithm, we need to tune the algorithm. The GUI for the software helps us in customizing the values for the algorithm. We can setup initial crossover and mutation rates and population size using the graphical frontend as shown in Figure 40. Similarly, number of hidden layers in neural network can be modified as shown in Figure 41.

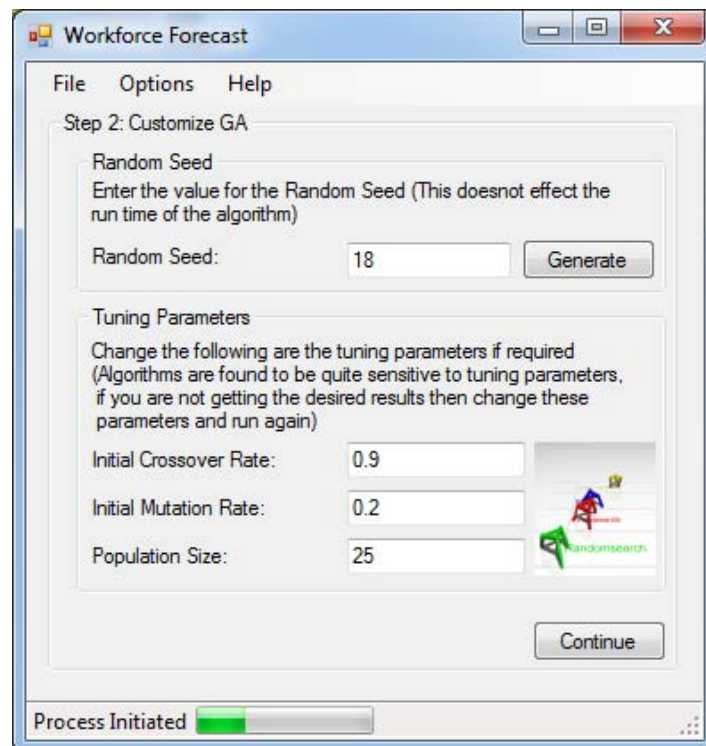


Figure 40: Customizing the GA Parameters

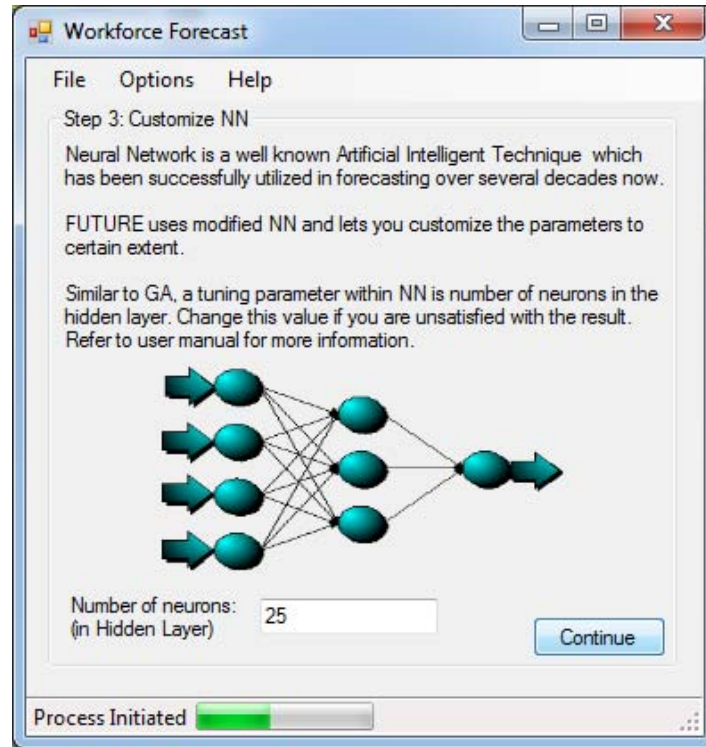


Figure 41: Customizing NN Parameters

Thus, an overall setting for the tuned performance of the algorithm for this problem is summarized in Table 8. Thereafter, SGA GONN is trained by the training dataset. Number of nodes in the input layer is 17, in hidden layer is 25 and in the output layer is 2. The tuned GA parameters were found to be as follows: crossover rate = 0.9 and the initial mutation rate is set as 0.2. The evaporation factor for the self guided ants is $r = 0.1$ and initial pheromone value is $t_0 = 0.9$. The weights for connections were randomly initialized between -5 to 5 for the tuned results.

Table 8: Parameter Settings

Attribute	Parameter Value
Initial Crossover Rate	0.9
Initial Mutation Rate	0.2
Population size	25
Number of neurons in:	
Input layer	17
Hidden Layer	25
Output Layer	2
Number of Ants	17
t_0	0.9
r	0.1

The error during the training is 0.95% and the error during the testing phase is 2.19% as compared to a Genetic Algorithm Optimized Neural Network (GANN) where the error in training was found to be 3.22% and error during the testing was found to be 4.68%. The comparative analysis with conventional Neural Network (backpropagation algorithm) has also been carried out. The training error in the conventional NN was found out to be 4.23% and the testing error was found out to be 6.79%. By utilizing the trained SGA GONN and the survey questionnaire for the year to be forecasted the forecasting year workforce came out to be 964 for workforce category 1 and 432 for workforce category 2. All these algorithms have been run for 10 times with different initial random seeds. The results presented above are the averaged results obtained for these 10 replications.

The performance is also compared with C^2 FDT (Sanjay, 2009) and the results obtained by SGA GONN significantly outperform it. Implementation of C^2 FDT algorithm on data for category 1 workforce results in a tree composed of 45 nodes. To carry out analysis of resulting tree its structure till 4th depth is shown in Figure 42. In the figure, inside nodes, first number represents node identity (i.e., node number) and numbers in parenthesis refer to total number of data points reside in that node. For any node, features like inconsistency index and node representative are determined by the number of data points and node center. Details of all these 15 nodes are provided in Table 9. In the table, node center, node representative, inconsistency index and number of data points in each node are shown. From the table it is evident that as C^2 FDT progresses in depth number of data points and inconsistency index is lower for children as compared to parent node. It clearly shows that as tree expands it tries to approach towards its structural optima.

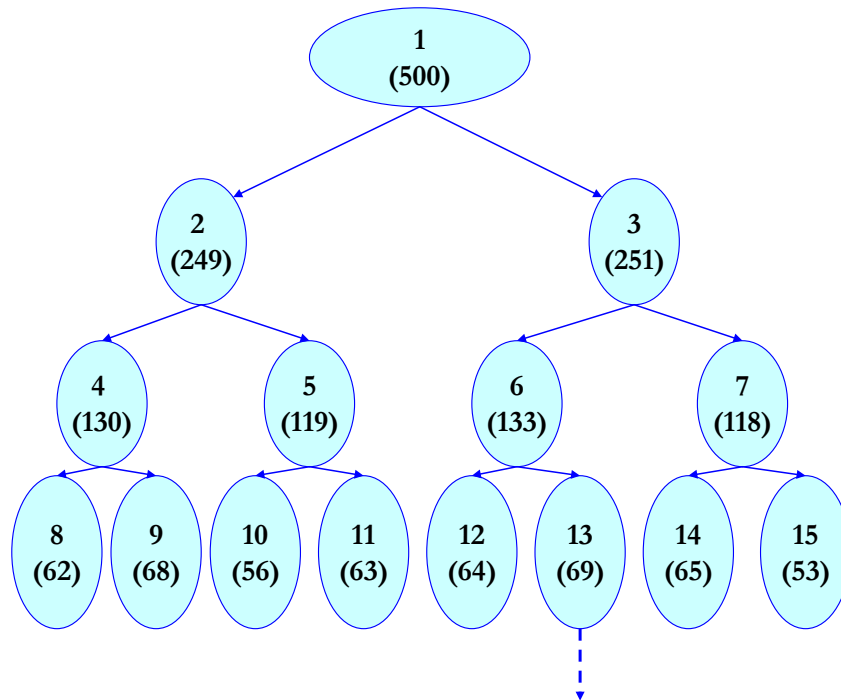
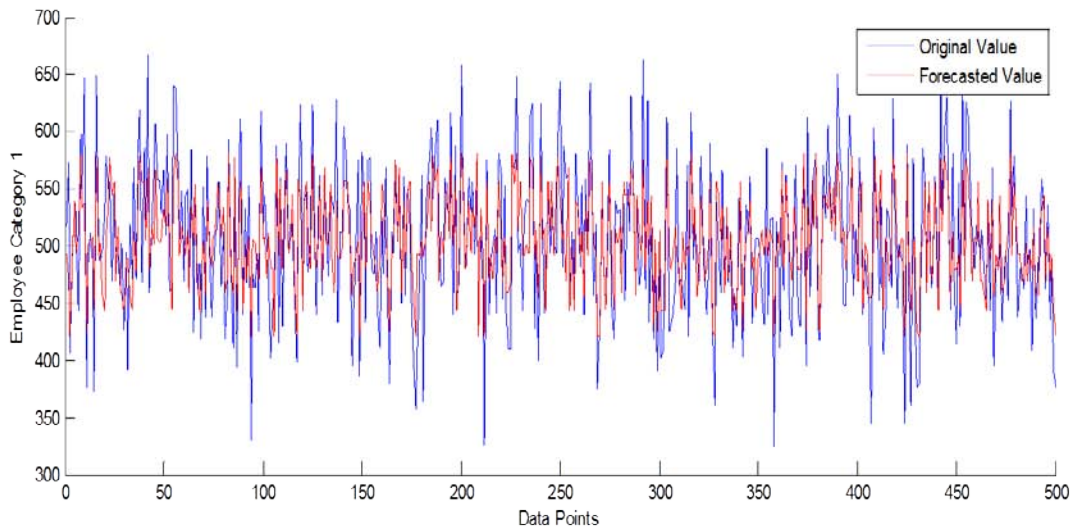


Figure 42: Structure of C^2 FDT for Workforce Category 1

Table 9: Description of C²FDT for Category 1 Employee Data Set

Node	Node's Center								Inconsistency	Node Representative	Total Data Points in node
2	0.5022	0.5034	0.5057	0.5070	0.4987	0.5026	0.5050	0.5030	479352.7497	526.3173	249
3	0.5022	0.5034	0.5057	0.5070	0.4987	0.5026	0.5050	0.5030	510241.6700	485.0757	251
4	0.4892	0.5087	0.5054	0.4997	0.6498	0.4522	0.4226	0.4985	133988.2523	485.5948	130
5	0.4895	0.5087	0.5050	0.4994	0.6526	0.4565	0.4236	0.4985	121955.8164	570.7860	119
6	0.5147	0.4980	0.5064	0.5144	0.3499	0.5537	0.5895	0.5073	166659.6577	522.8858	133
7	0.5153	0.4982	0.5062	0.5144	0.3448	0.5473	0.5830	0.5075	145800.2134	442.5349	118
8	0.4749	0.5096	0.5199	0.5094	0.5662	0.3139	0.3997	0.4983	64572.9402	483.6774	62
9	0.4749	0.5096	0.5199	0.5094	0.5662	0.3139	0.3997	0.4983	68420.6279	487.4706	68
10	0.5008	0.5042	0.4887	0.4772	0.7226	0.6387	0.3660	0.4998	72582.6004	559.8282	56
11	0.5104	0.5116	0.4895	0.4994	0.7646	0.5783	0.5278	0.4975	64890.4207	580.4413	63
12	0.5031	0.4912	0.5101	0.5146	0.4399	0.5934	0.7305	0.5006	75154.3671	504.6702	64
13	0.5247	0.4986	0.5108	0.5142	0.4188	0.7082	0.6423	0.5072	93250.5284	539.2443	69
14	0.5153	0.5057	0.5056	0.5231	0.2548	0.4817	0.4326	0.5085	62150.0778	464.3791	65
15	0.5166	0.4983	0.4973	0.5063	0.2567	0.3903	0.5106	0.5136	60526.2081	415.8004	53

After developing the C²FDT for workforce category 1, it is tested with 500 data points. The forecasting capability of C²FDT for workforce category 1 is shown in Figure 43. The training error in C²FDT is 7.87% and the testing error is 10.91%.

**Figure 43: Forecasting for Employee Category 1 Training Data Set by C²FDT**

In order to investigate the dependency of forecasting accuracy on data size, various test instances are generated. For all of these instances C²FDT is developed and tested. Results obtained after conducting rigorous experimentation are shown in Figure 44. From the figure it is evident that, in general, average percentage (%) error in forecasting was decreased significantly up to data set of 1000 points. Afterwards, with the increase in size of data there is no significant reduction in average percentage error, while, increase in huge computational time was witnessed.

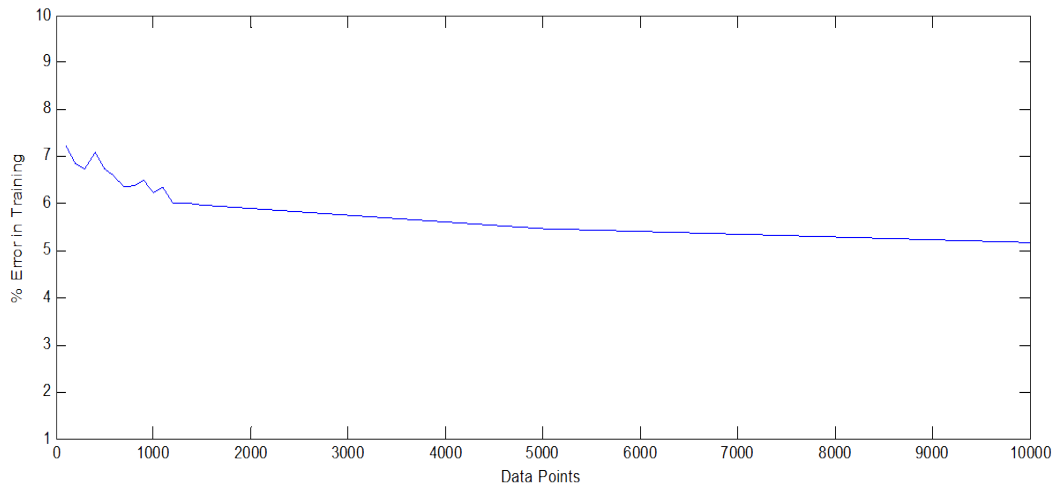


Figure 44: Average Percent Error vs. Size of Data Set in Training of C²FDT for Employee Category 1

Figure 45 shows the convergence trend for the SGA GONN and its comparison with GANN and C²FDT. It can be seen that SGA GONN clearly outperforms the GANN and C²FDT algorithm.

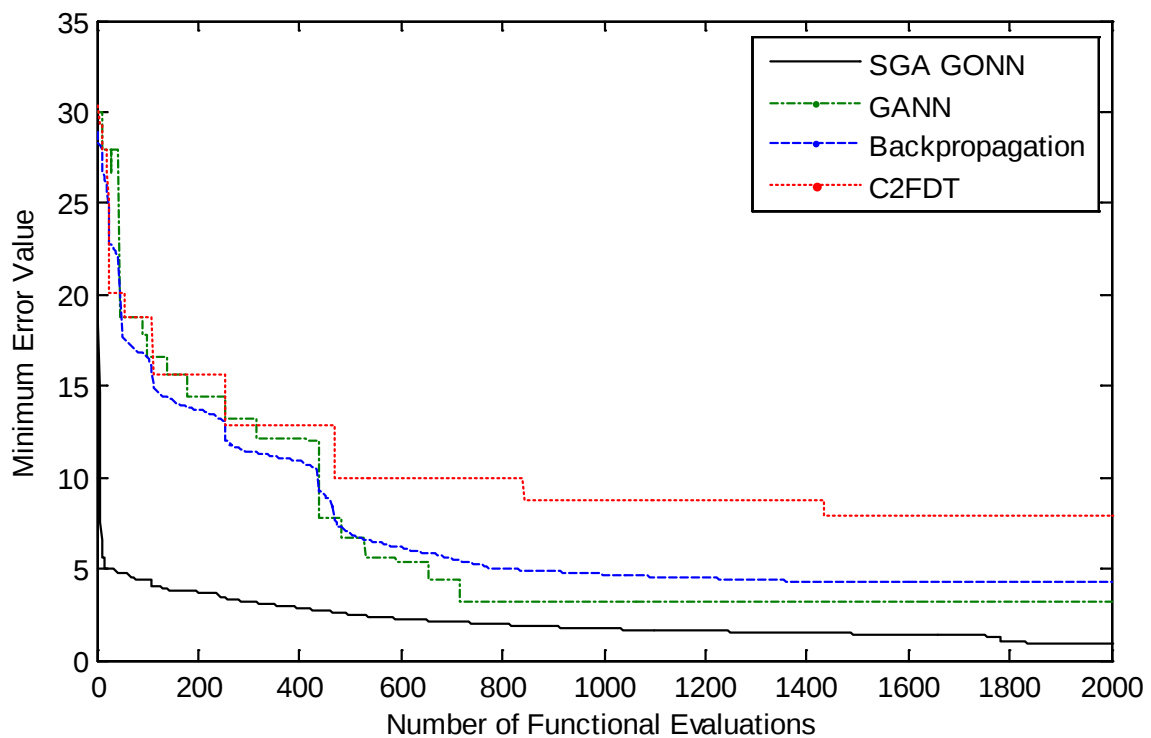


Figure 45: Comparative Analysis of the Convergence Trends for Three Algorithms

Figure 46 shows the comparison between the actual dataset and the adapted results from the algorithm for workforce category 1. From the same figure it is visible that the results produced by SGA GONN almost superimpose the original dataset.

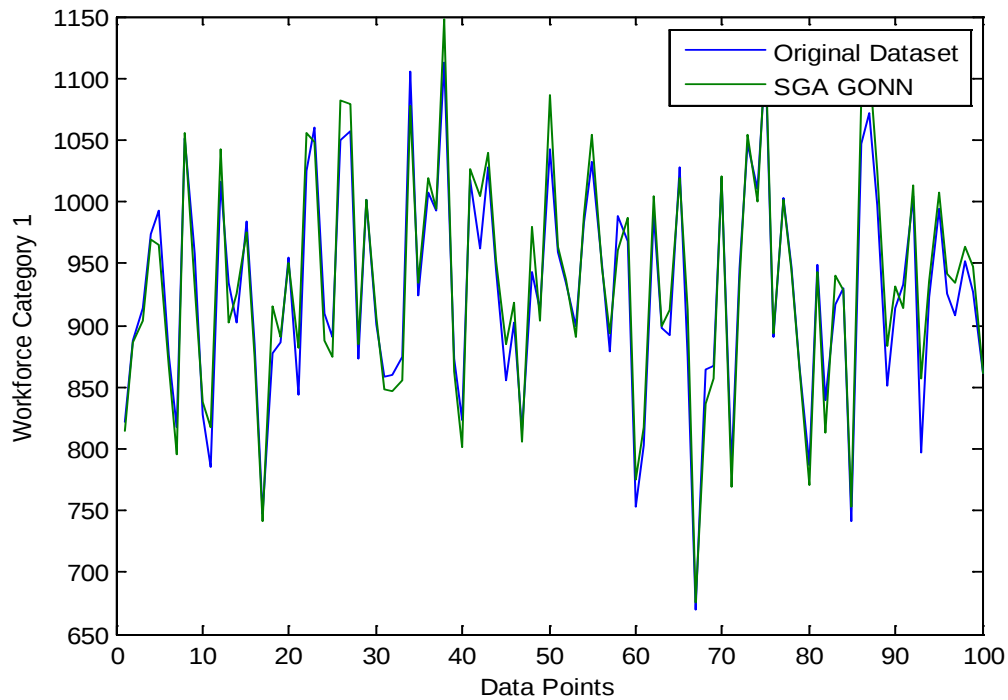


Figure 46: SGA GONN for the Training Data Set (Workforce Category 1)

Figure 47 shows the comparison between the actual dataset and the adapted results from the algorithm for workforce category 2. Again, from the same figure it is visible that the results produced by SGA GONN almost superimpose the original dataset.

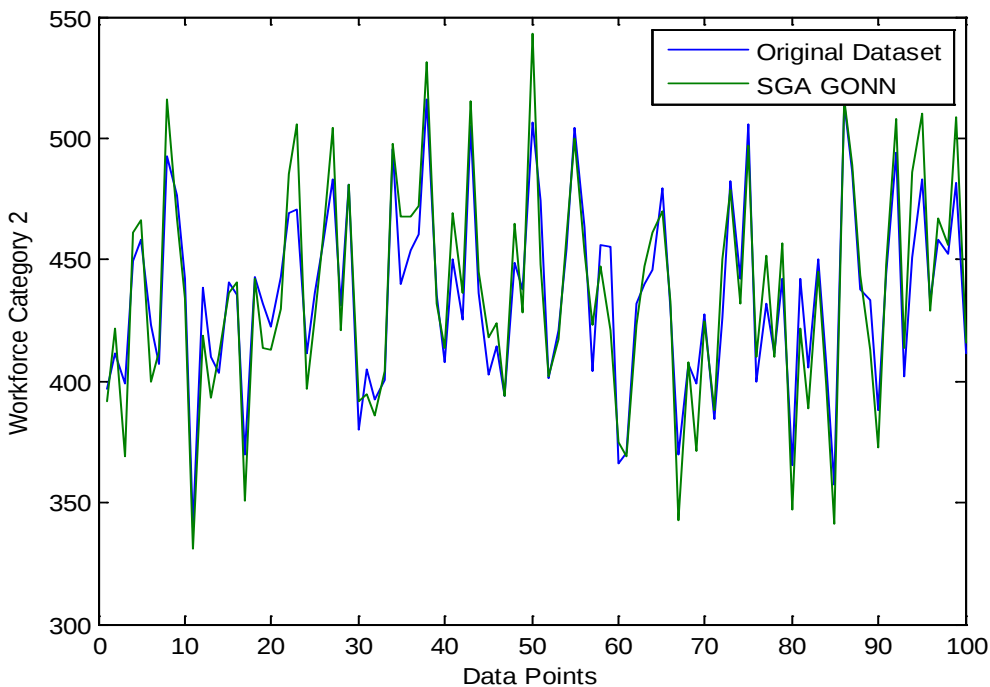


Figure 47: SGA GONN for the Training Data Set (Workforce Category 2)

As already seen in Figure 45, SGA GONN is capable of reaching much lower error values as compared to the other two algorithms. Same results can also be viewed in the figures above where SGA GONN is able to get closer to the data points in the testing dataset. Similarly, in Figure 48 and Figure 49, SGA GONN clearly outperforms GANN and C²FDT.

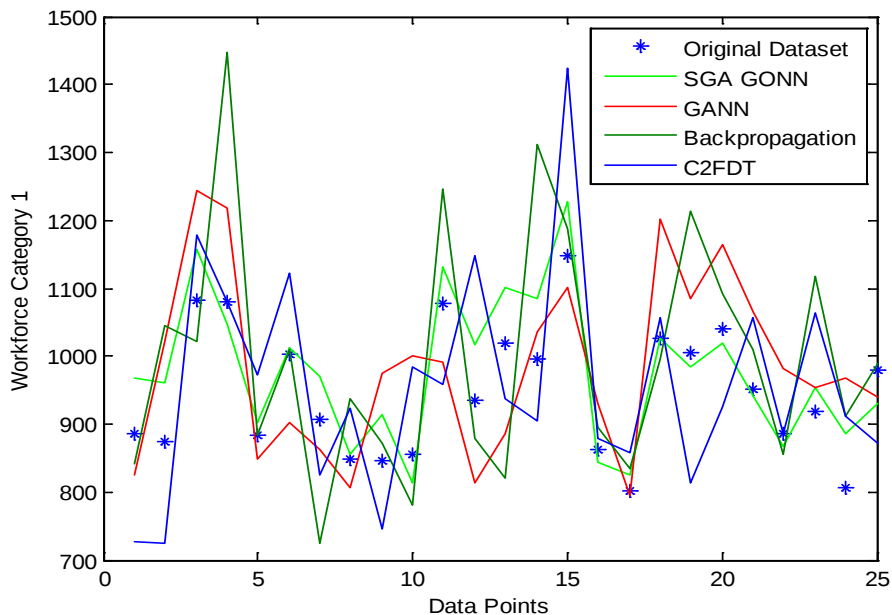


Figure 48: Comparative Analysis for Three Algorithms over Workforce Category 1 (Testing Dataset)

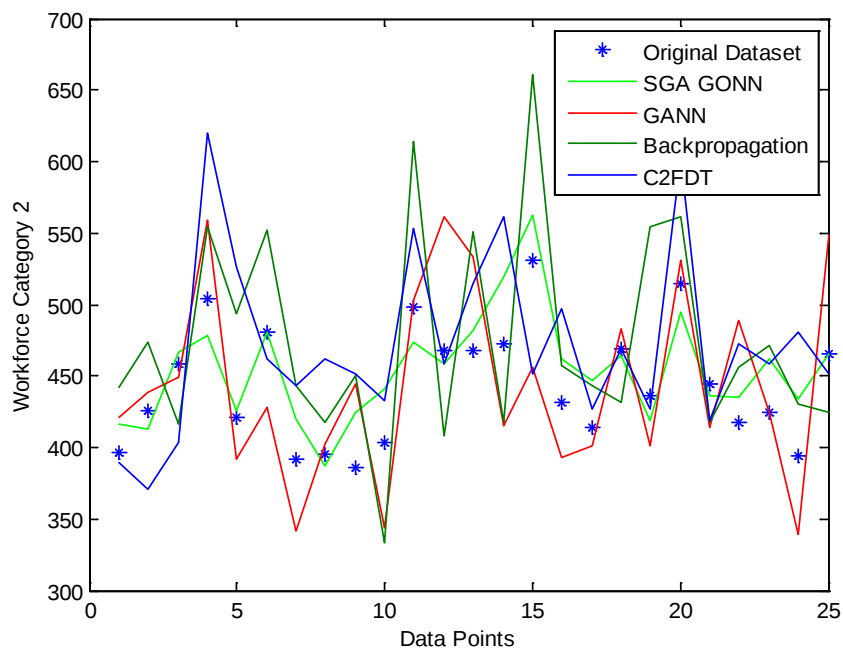


Figure 49: Comparative Analysis for Three Algorithms over Workforce Category 2 (Testing Dataset)

13.5 Limitations of this Research

This research presents an innovative approach to deal with the workforce forecasting problem. However, in its present form, there are several limitations of this research which are discussed as below:

1. **Data Collection:** This research relies on the past data for learning the existing pattern in the workforce forecasting parameters and the net workforce. Collecting 50 data points or more for past 20 years is not practical approach, therefore, the software cannot be used to predict future workforce immediately. However, it is an iterative process and the software collects data over the period of time and adapts to the changes accordingly.
2. **Time Series and Sudden Unexpected Changes:** The structure of the SGA GONN is designed in such a way that it may not be able to consider sudden unexpected changes or time-series problems like lagging. Besides, the relationship between the input parameters and the workforce might change significantly over the period of time. It should be noted that the current version of the software gives equal preference to all the past relationships.
3. **Workforce Mix:** The current software only forecasts 2 categories of workforce, however, the same methodology can be extended to forecast more categories of workforce.
4. **Workforce Deployment:** The results reported by the software suggest the total number of the final workforce. However, it does not provide any guideline to deal with the changes in the workforce over a period of time. These decisions will be made by the higher level management (HR committee).

14.0 CONCLUSIONS

This project studies workforce forecasting in two aspects: (1) an extensive review of the existing methodologies and techniques and (2) an effort to develop a decision support system with models and software programming. In this report (i.e., Part II of this research) a forecasting decision support system has been developed for forecasting the future workforce of an organization. The workforce planning methodology proposed in this research:

1. Develops an iterative and incremental process for the workforce forecasting,
2. Integrates the strategic business goals with workforce planning process,
3. Suggests a longer planning horizon for an accurate forecasts,
4. Implements a broad domain of workforce planning activities, from environmental scanning to strategic reviews,
5. Reduces the future uncertainty by incorporating the skill's gap analysis, and
6. Provides a competency framework which supports coherent analysis of several factors involved.

Furthermore, by utilizing various soft-computing techniques, this research proposes a new data mining technique named as Self Guided Ant-based Genetically-Optimized-Neural-Network (SGA GONN). The proposed SGA GONN has its roots in the traditional Genetic Algorithm (GA), Neural Network (NN), Algorithm of Self Guided Ants (ASGA) and Fuzzy Logic. The software is coded and compiled in MATLAB and the graphical user interface is developed in Visual C#.

Thereafter a case study of a hypothetical organization has been simulated where the goal is to forecast the size and mix of workforce for 5 years. After simulating the survey by using MATLAB simulation, 125 data points are generated to test the model. Each one of the data points contains 17 parameter values and information of 2 output workforce categories which has been generated by simulating a random function. Thereafter, the proposed SGA GONN is utilized to predict the future workforce of the organization. The algorithm performance is compared with a few other data mining techniques available in literature e.g., C²FDT and GANN where SGA GONN consistently produces superior results on the problem at hand.

This research provides several unique contributions to the area of workforce forecasting and data mining soft computing.

1. This research presents a more comprehensive set of questionnaire with 52 questions and 17 parameters to capture the human resource information within an organization in an extensive manner.
2. This research presents an improved intelligent decision support system for workforce forecasting. After conducting several tests on various artificial intelligence techniques, it was found that neural network based approaches significantly outperform the decision tree approaches. The proposed SGA GONN is a robust solution methodology that outperforms other existing solution methodologies such as GANN or C²FDT.

3. A GUI has been developed in Visual C# which provides the capability of customizing the algorithm and also helps in maintaining the database. Thus, users with little knowledge of GA or NN will still be able to operate the software, maintain the database and execute the algorithm by simply following the instructions with suggested default settings.

REFERENCES

- Ada, G.L. and Nossal, G. (1987). The clonal selection theory. *Scientific American*, 257(2), 50-57.
- Anderson, E.J. and Ferris, M.C. (1989). *A genetic algorithm for the assembly line balancing problem* (Technical Report). University of Wisconsin-Madison, Computer Sciences Department, Madison, WI.
- Bezdek, J.C. (1981). *Pattern recognition with fuzzy objective functions*. New York: Plenum.
- Brown, M. and Harris, C. (1994). *Neural fuzzy adaptive modeling and control*. Englewood Cliffs, Prentice-Hall, NJ.
- Cleveland, G.A. and Smith, S.F. (1989). Using genetic algorithms to schedule flow shop releases. In *Proceedings of 3rd International Conference on Genetic Algorithm and Their Applications* (pp. 160-169). Palo Alto, CA.
- Cotton, A. (2007). *Seven steps of effective workforce planning*. Human Capital Management-Series (<http://www03.ibm.com/industries/government/us/detail/resource/X448444G98379A82.html>).
- Cox, L.A., Davis, L., and Qiu, Y. (1991). Dynamic anticipatory routing in circuit-switched telecommunications networks. In L. Davis (Ed.), *Handbook of genetic algorithms*. New York: Van Nostrand Reinhold.
- Davis, L. (1985). Applying adaptive algorithms to epistatic domains. In *Proceedings of the 9th International Joint Conference on Artificial Intelligence* (Vol. 1, pp. 162-164). August, Los Angeles, CA.
- De Castro, L.N. and Timmis, J.I. (2002). *Artificial immune systems: A new computational intelligence approach*. London: Springer-Verlag.
- Deb, K. (2004). *Single- and multi-objective optimization using evolutionary computation* (Technical Report No. 2004002). Kanpur, India: Kanpur Genetic Algorithms Laboratory (KanGAL).
- Dobra, A. and Gehrke, J. (2002). SECRET: A Scalable Linear Regression Tree Algorithm. In *The 8th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*. July 23-26, Edmonton, Alberta, Canada.
- Dorigo, M. and Gambardella, L.M. (1997). Ant colonies for the traveling salesman problem. *BioSystems*, 43, 73-81.
- Dorigo, M. and Stützle, T. (2004). *Ant colony optimization*. Cambridge, MA: MIT Press.
- Dorigo, M. (2004). *Ant colony optimization and swarm intelligence, 4th International Workshop, ANTS 2004*. September 5-8, Brussels, Belgium, Berlin: Springer.
- Dorigo, M. (2006). *Ant colony optimization and swarm intelligence, 5th International Workshop, ANTS 2006*. September 4-7, Brussels, Belgium, Berlin: Springer.
- Dorigo, M. (2008). *Ant colony optimization and swarm intelligence, 6th International Workshop, ANTS 2008*, September 22-24, Brussels, Belgium, Berlin: Springer.
- Eggermont, J. (2002). Evolving fuzzy decision trees with genetic programming and clustering. *Proceeding of the 5th European Conference on Genetic Programming (EuroGP'02)*. *Lecture Notes in Computer Science*, 2278, 204-237.
- Fourman, M.P. (1985). Compaction of symbolic layout using genetic algorithms. In *Proceedings of the First International Conference on Genetic Algorithms and their Applications* (pp.141-155). July, Pittsburgh, PA.

- Franks, N.R. and Richardson, T. (2006). Teaching in tandem-running ants. *Nature*, 439, 153-154.
- Franks, N.R., Dornhaus A., Fitzsimmons J. P., and Stevens M. (2003). Speed versus accuracy in collective decision making. *Proceedings of the Royal Society B: Biological Sciences*, 270, 2457-2463.
- Gen, M., Tsujimura, Y., and Li Y. (1996). Fuzzy assembly line balancing using genetic algorithms. *Computers and Industrial Engineering*, 31(3/4), 631-634.
- Goldberg, D.E. (1987a). Computer-aided pipeline operation using genetic algorithms and rule learning. Part I: Genetic algorithms in pipeline optimization. *Engineering with Computers*, 3(1), 35-45.
- Goldberg, D.E. (1987b). Computer-aided pipeline operation using genetic algorithms and rule learning. Part II: Rule learning control of a pipeline under normal and abnormal conditions. *Engineering with Computers*, 3(1), 47-58.
- Goldberg, D.E., Korb, B., and Deb, K. (1989). Messy genetic algorithm: Motivation analysis and first results. *Complex Systems*, 3, 493-530.
- Grefenstette, J.J., Gopal, R., Rormatia, B., and Vangucht, D. (1985). Genetic algorithms for traveling salesman problem. In *Proceedings of the First International Conference on Genetic Algorithms and their Applications*. July, Pittsburgh, PA.
- Han, J. and Kamber, M. (2001). *Data mining concepts and techniques*. San Francisco, CA: Morgan Kaufmann.
- Hopfield, J.J. and Tank, D.W. (1985). Collective computation with continuous variables. In E. Bienstock, F. Fogelman, and G. Weisbuch (Eds.), *Distorted system and biological organization*. Berlin: Springer Verlag.
- Hopfield, J.J. (1982). Neural networks and physical systems with emergent collective computational abilities. *Proceedings of the National Academy of Science*, 79, 2554-2558.
- International Personnel Management Association (IPMA-HR) (2000). *Workforce planning resource guide for public sector human resource professionals*. U.S. Department of Homeland Security, Alexandria, VA.
- Jerne, N.K. (1974). Towards a network theory of the immune system. *Annals of Immunology*, 125C(1-2), 373-389.
- Keel, J. (2006). *Workforce planning guide*. State Auditor's Office, Edition SAO Report No.06-704. (<http://sao.hr.state.tx.us/Workforce/06-704.pdf>).
- Lee, C.Y., Piramuthu, S., and Tsai, Y.K. (1997). Job shop scheduling with genetic algorithm and machine learning. *International Journal of Production Research*, 35(4), 1171-1191.
- Maulik, U. and Bandyopadhyay, S. (2003). Fuzzy partitioning using a real-coded variable length genetic algorithm for pixel classification. *IEEE Transactions on Geosciences and Remote Sensing*, 41(5), 1075-1081.
- Mendes, R.R.F., Voznika, F.B., Freitas, A.A., and Nievola, J.C. (2001). Discovering fuzzy classification rules with genetic programming and co-evolution. In 5th European Conference on Principles of Data Mining and Knowledge Discovery. *Lecture Notes in Computer Science*, 2168, 314-325.
- Mierswa, I. (2005). Incorporating fuzzy knowledge into fitness: multiobjective evolutionary 3D design of process plants. In *Proceedings of Genetic and Evolutionary Computation Conference* (pp. 1985-1992). June 25-29, Washington, DC.
- Nossal, G.J. (1994). Negative selection of lymphocytes. *Cell*, 76, 229-239.

- Pedrycz, W. and Sosnowski, Z.A. (2000). Designing decision trees with the use of fuzzy granulation. *IEEE Transactions on Systems, Man, and Cybernetics-Part A*, 30(2), 151-159.
- Pedrycz, W. and Sosnowski, Z.A. (2005). C-fuzzy decision trees. *IEEE Transactions on Systems, Man, and Cybernetics-Part C*, 35(4), 498-511.
- Pham, D.T. and Karaboga, D. (2000). *Intelligent optimization techniques: Genetic algorithms, tabu search, simulated annealing and neural networks*. New York: Springer.
- Quagliarella, D. (1998). *Genetic algorithms and evolution strategy in engineering and computer science: Recent advances and industrial applications*. Chichester, NY: John Wiley & Sons.
- Ross, T.J. (1997). *Fuzzy logic with engineering applications*. New York, NY: McGraw Hill.
- Russel, S. and Norvig, P. (1995). *Artificial intelligence: A modern approach*. Englewood Cliffs, NJ: Prentice-Hall.
- Sanchez, E., Shibata, T., and Zadeh, L.A. (1997). *Genetic algorithms and fuzzy logic systems: Soft computing perspectives*. River Edge, NJ: World Scientific Publications.
- Shukla, S.K. and Tiwari, M.K. (2009). Soft decision trees: A genetically optimized cluster oriented approach. *Expert Systems with Applications*, 36(1), 551-563.
- Shukla, S.K. (2009). *Decision models and artificial intelligence in supporting workforce forecasting and planning* (Master's Thesis). University of Texas at San Antonio, TX.
- Shukla, S.K., Tripathi, M., Kuriger, G., Wan, H., and Chen, F.F. (2009). Clonal C-fuzzy decision tree (C²FDT) for workforce deployment. In *2009 Annual Industrial Engineering Research Conference* (pp. 2128-2133). May 30-June 3, Miami, FL.
- Spears, W. (1997). Recombination parameters. In T. Back, D.B. Fogel, and Z. Michalewicz (Eds.), *The handbook of evolutionary computation* (pp. E1.3:1-E1.3:13). Philadelphia, PA: IOP Publishing and Oxford University Press.
- Thompson, J.R. and Mastracci, S.H. (2005). *The blended workforce: Maximizing agility through nonstandard work arrangements*. IBM Center for the Business of Government, Washington, DC.
- Tripathi, M., Agrawal, S., and Tiwari, M.K. (2008). Disassembly sequencing problem: Resolving the complexity by random search techniques. In S. M. Gupta and A. J. D. Lambert (Eds.), *Environment conscious manufacturing* (pp. 331-362). Boca Raton, FL: CRC Press.
- Tripathi, M., Shukla, S.K., Kuriger, G., Wan, H., Chen, F.F., and Riehl, B.D. (2010). Forecasting decision support system: A self-guided ant based genetically-optimized-neural-network approach. In *2010 Annual Industrial Engineering Research Conference*. June 5-9, Cancun, Mexico.
- Tseng, L.Y. and Yang, S.B. (2001). A genetic approach to the automatic clustering problem. *Pattern Recognition*, 34, 415-424.
- Weber, R. (1992). Fuzzy ID3: A class of methods for automatic knowledge acquisition. In *Proceedings of the 2nd International Conference on Fuzzy Logic Neural Networks* (pp. 265-268). Iizuka, Japan.
- Womack, J.P. and Jones, D.T. (2003). *Lean thinking: Banish waste and create wealth in your corporation* (2nd ed.). New York, NY: Simon & Schuster.
- Yao, X. (1999). Evolving artificial networks. *Proceedings of the IEEE*, 87(9), 1423-1447.
- Yildiz, O.T. and Alpaydin, E. (2001). Omnivariate decision trees. *IEEE Transactions on Neural Networks*, 12(6), 1539-1546.

LIST OF ACRONYMS

ACA	Adaptive Conjoint Analysis
ACO	Ant-Colony Optimization
AIS`	Artificial Immune System
ANN	Artificial Neural Networks
AR	Auto-Regressive
ARARMA	Long-term and Short-term Auto-Regressive Moving Average
ARIMA	Auto-Regressive, Integrated, Moving-Average
ASGA	Algorithm of Self Guided Ants
BVAR	Bayesian Vector Autoregressive
CAA	Computer-Assisted Asynchronous
CART	Classification and Regression Trees
C ² FDT	Clonal C-Fuzzy Decision Tree
CFDT	Cluster Based Fuzzy Decision Trees
CIO	Chief Information Officer
CIVFORS	Civilian Forecasting System
CPDF	Civilian Personnel Data File
CSB	Central Statistical Bureau
CTAR	Continuous-Time Threshold Autoregressive
CUSUM	Cumulative Sums
CX	Cycle Crossover
DOD	Department of Defense
DON	Department of Navy
DOT	Department of Transportation
DT	Decision Trees
ECM	Error-Correction Model
EU	European Union
FCAR	Functional Coefficient Autoregressive
FCM	Fuzzy C-Mean Clustering
FDSS	Forecasting Decision Support System
FLC	Fuzzy Logic Controller
GA	Genetic Algorithm
GANN	Genetic Algorithm Optimized Neural Network
GCSDT	Genetically Optimized Cluster Oriented Soft Decision Trees
GDP	Group Decision Program
GIS	Geographic Information Systems
GUI	Graphical User Interface
HR	Human Resource
ID3	Iterative Dichotomiser 3
ID4	Incremental Induction of Decision Tree
IM/IT	Information Management/Information Technology)
INGT	Improved Nominal Group Technique
IPMA	International Personnel Management Association
MA	Moving-Average
MCR	MATLAB Component Runtime

MISO	Multiple-Input-Single-Output
NCAS	Non-Computer-Assisted Synchronous
NGT	Nominal Group Technique
PDP	Parallel Distributed Process
PMX	Partially Matched Crossover
PPX	Precedence Preservative Crossover
RBF	Rule-based Forecasting
ROI	Return on Investment
SES	Single Exponential Smoothing
SETAR	Self-Exciting Threshold Autoregressive
SGA GONN	Self Guided Ant-based Genetically-Optimized-Neural-Network
STAR	Smooth Transition Autoregressive
SWOT	Strengths, Weaknesses, Opportunities and Threats
TOPSIS	Technique for Order Preference by Similarity to Ideal Solution
TSP	Traveling Salesman Problem
VAR	Vector Auto-regression
VARIMA	Vector ARIMA