

Process Integration and Optimisation of the Wits Microbrewery Plant Using Artificial Neural Networks



EZEKIEL MAKWADINKGA MADIGOE

Process integration and optimisation of the Wits microbrewery plant using Artificial Neural Networks

EZEKIEL MAKWADINKGA MADIGOE

A thesis submitted to the Faculty of Engineering and the Built Environment,
University of the Witwatersrand, Johannesburg, in fulfilment of the requirements
for the degree of Doctor of Philosophy.

Johannesburg, 2013

DECLARATION

I declare that this thesis is my own unaided work unless otherwise stated. It is being submitted for the Degree of Doctor of Philosophy in the University of the Witwatersrand, Johannesburg. It has not been submitted for any degree or examination in any other University.

.....

Ezekiel M. Madigoe

02 October 2013

ABSTRACT

Beer brewing is one of the most widely practiced processes in the world. The brewing technology is associated with huge business opportunities because of high demand of the final product. Due to worldwide competition and increasing costs of resources and input costs, businesses have to consider innovative ways of brewing beer optimally. Optimisation of the brewing process is a challenging task that relies on understanding the complex effects of processing inputs on productivity and quality outputs, whose relationship exhibits non-linear behaviour. In this study, various optimisation methods aimed at reducing feed and energy costs are introduced. The proposed methods use Response Surface Methodology and neural networks for optimisation of significant inputs. Process integration is used for increased energy recovery and to minimize waste energy; and the prospects of using fed-batch fermentation in brewing are evaluated.

There is a need for more advanced optimisation tools that can handle complex behavior of biological processes such as brewing process. The importance of this study is reflected in the novel use of Artificial Neural Networks with Genetic Algorithms (ANN) as an optimisation tool for biological systems. Because of its great potential to circumvent the limitations associated with the traditional optimisation methods, it can serve as a basis of discernment in improving the brewing process.

Increasing energy costs as well as inefficiencies optimisation of batch brewing processes as compared to continuous processes make it necessary to save energy through other strategies such as process integration; and possible switch to fed-

batch fermentation. The thesis makes significant contributions to current debates on switching to fed-batch fermentation in the brewing industry. It also extends the existing focus on batch process integration for sustainable industries.

Artificial Neural Network models proved superior and more accurate than Response Surface Methodology model when tested against various experimental results. The optimal values of temperature, sugar concentration and yeast concentration were found to be 19.0°C, 163.9 g/l and 0.42 g/l, respectively. The optimal values from the Genetic Algorithm were then inserted as inputs to the neural network to yield an optimal ethanol concentration of 63.49 g/l.

Fed-batch fermentations produced higher ethanol yields than batch, although the differences were not major to justify the switch to fed-batch fermentation. After heat integration of process units was deployed to the Wits Microbrewery plant, heat loss and energy demand were reduced by 19.5% and 23.0%, respectively. It was observed that besides time and temperature constraints, energy load mismatch is one of the critical variables that restrict maximum energy recovery in batch process integration.

It was thus concluded that ANN is suitable for prediction and optimisation of complex and non-linear processes like brewing fermentation. It was also concluded that opportunities exist for energy savings and costs reduction through implementation of heat integration strategies and use of fed-batch fermentation in brewing process.

PUBLICATIONS AND PRESENTATIONS

Journal Publications

1. **Madigoe, E. M.**, Iyuke, S. E., Simate, G., Lefofana, K, Mthembu, S. (2012). Formulation and optimisation of empirical models using Response Surface Methodology co-relating medium concentration, temperature, and yeast concentration to ethanol production during fermentation, *South African Journal of Chemical Engineering*, vol. 16, no 1, 76-10.

Peer Reviewed Conference Proceedings

1. **Madigoe, E. M.**, Iyuke S. E. (2012). Comparison of batch and fed-batch fermentation at Wits Microbrewery plant, 2nd International Conference on Challenges in Biotechnology and Food Technology, Annamalai University of Tamil Nadu State, India, , 9-10 January, 2012.
2. **Madigoe, E. M.**, Iyuke S. E (2012). Basic comparison of batch and fed-batch fermentation at Wits Microbrewery plant, South African Institute of Chemical Engineers (SAICHE) Conference, Champagne Sports Resort, Drakensberg, South Africa, 16-19 September, 2012.

DEDICATION

Dedicated to my mother, Ramakgahlela Madigoe...Hunadi...! And my father, Mogoboya Madigoe...Nape...! Who have supported me all the way in everything I do.

To my family and friends who have been great source of motivation and inspiration.

Finally, this thesis is dedicated to all those who believe in the richness of learning.

ACKNOWLEDGEMENTS

I am grateful to my supervisor, Professor Sunny Iyuke for dedicating most of his time and affording the opportunity into a successful Doctorate research. I am thankful for his guidance and encouragement to develop me into an independent researcher. Special thanks to all the students I had pleasure in supervising in the final year of their studies, you are much appreciated. Special thanks also are also due to all at Biochemical Engineering and Nano-technology Research Group.

I acknowledge the support of SAB Miller Alrode, South Africa specifically, Anton Erasmus, for the supply of raw materials.

Special thanks to my family and friends for their consistent support throughout this work and for believing in me. I would like to acknowledge the financial support of Sasol, National Research Foundation, FoodBev SETA and Wits University. This publication was made possible (in part) by a grant from the Carnegie Corporation of New York. The statements made and views expressed are, however, solely the responsibility of the author.

Through Him ... With Him... In Him ...

Thank you for the Blessings, the Almighty Lord

Contents

DECLARATION	iii
ABSTRACT	iv
PUBLICATIONS AND PRESENTATIONS	vi
DEDICATION	vii
ACKNOWLEDGEMENTS.....	viii
LIST OF SYMBOLS	xvi
1. INTRODUCTION	2
1.1 Introduction	2
1.2 Research problem and contribution to knowledge	6
1.3 Aim and objectives.....	7
1. 4 Research Method.....	8
1.5 Thesis Outline.....	8
2. LITERATURE REVIEW.....	14
2.1 General Introduction.....	14
2.1.1 Brewing Industry	15
2.2 Beer Brewing Process.....	17
2.2.1 Background	17
2.2.2 Malting	17
2.2.3 Milling	18
2.2.4 Mashing.....	19
2.2.5 Lautering	19
2.2.6 Boiling.....	21
2.2.7 Fermentation	21
2.3 Challenges in optimisation of brewing fermentation	22
2.3.1 Background	22
2.3.2 Different types of fermentation processes.....	24
2.3.3 Batch and Fed-Batch Fermentation	25
2.3.4 Fermentation Modelling	31
2.4 Response Surface Methodology	33
2.4.1 Modelling Response Surface Methodology (RSM)	33
2.4.2 Central Composite Rotatable Design	36

2.4.4 Determining the Coordinates of the Stationary Point	39
2.5 Artificial Neural Networks (ANN)	41
2.5.1 Development of Neural Networks	42
2.5.2 Architecture of the Artificial Neural Network.....	44
2.5.3 Application of Artificial Neural Networks	48
2.5.4 Importance of Artificial Neural Networks.....	50
2.6 Genetic Algorithms	52
2.6.1 Why Genetic Algorithms?	53
2.7 Process Integration	55
2.7.1 Process integration for batch processes.....	55
2.7.2 Sequential approach	59
2.7.3 Simultaneous framework.....	61
2.7.4 Pinch analysis/ technology.....	63
2.7.5 Time consideration methods	64
3. MATERIALS AND METHODS.....	71
3.1 Introduction	71
3.2 Materials	71
3.2.1 Wort preparation	71
3.2.2 Reagents.....	74
3.2.3 Yeast Preparation.....	74
3.3 Experimental Methods.....	74
3.3.1 Development of Artificial Neural Network model for brewing fermentation system.	74
3.3.2 Reliability of Response Surface Methodology for analysis and optimisation compared to Artificial Neural Networks.	75
3.3.3 Comparison of fed-batch fermentation vs. batch fermentation	77
3.3.4 Integration of the brewing process.....	78
3.4 Analytical Techniques	78
3.4.1 Validation of the ANN model.....	78
3.4.2 Yeast viability	80
3.4.3 Fermentable sugars.....	80
3.4.4 Specific gravity	82
3.4.5 Ethanol content.....	83

3.5 Data Analysis	84
3.6 Summary	84
4. OPTIMISATION USING RESPONSE SURFACE METHODOLOGY.....	86
4.1 Introduction	86
4.2 Materials and methods	87
4.2.1 Experimental Design for RSM and CCD.....	87
4.3 Results and discussion	92
4.3.1 Derivation of fitted model	92
4.3.2 Checking the adequacy of the developed model	94
4.3.3 Response surfaces.....	95
4.3.4 Combined effects of temperature and yeast concentration	96
4.3.5 Combined effects of Temperature and sugar concentration	98
4.3.6 Combined effects of sugar and yeast concentration	100
4.3.6 Determination of Optimum Conditions	102
4.15 Summary and conclusions	105
5. OPTIMISATION USING ARTIFICIAL NEURAL NETWORKS.....	109
5.1 Introduction	109
5.2 Materials and methods	110
5.2.1 Modelling the Artificial Neural Network.....	110
5.2.2 Data for artificial neural network (ANN).....	110
5.2.3 Neural network training, validation and testing	112
5.2.4 Choosing the Artificial Neural Network Architecture	112
5.2.5 Determination of optimal values	114
5.2.5 Checking the adequacy of the developed ANN model	115
5.3 Results and Discussion	118
5.3.1 Artificial Neural Network (ANN) model	118
5.3.2 Performance of the Genetic Algorithm.....	126
5.3.3 Genetic Algorithm aspects during optimisation	131
5.6 Supporting GA plots for an optimum Ethanol concentration.....	134
5.4 Summary and conclusions	135
6. INTEGRATION AND OPERATION OF WITS MICRO-BREWERY.....	138
6.1 Comparison of batch and fed-batch fermentation.....	138
6.1.1 Introduction	138

6.1.2 Materials and Methods.....	139
6.1.3 Yield and productivity of batch and fed-batch fermentation at various temperatures.	139
6.1.4 Comparison between experimental and theoretical data correlations.	145
6.1.5 Summary and conclusions	147
6.2 Integration of the Wits Micro-brewery plant	148
6.2.1 Introduction	148
6.2.2 Materials and methods.....	149
6.2.3 Result and Discussion.....	152
6.2.4 Summary and conclusions	157
7. CONCLUSIONS AND RECOMMENDATIONS.....	160
7.1 Conclusions	160
7.1.1 Introduction	160
7.1.2 Optimisation using Response Surface Methodology	161
7.1.3 Optimisation using Artificial Neural Networks	162
7.1.4 Comparison of batch and fed-batch fermentation.....	163
7.1.5 Integration of the brewing plant.....	164
7.1.6 Potential application in industry.....	164
7.2 Recommendations	165
REFERENCES.....	167
APPENDIX A: RSM MODEL VALIDATION	189
APPENDIX B: RSM SAMPLE CALCULATIONS	199
APPENDIX C: HIGH PERFORMANCE LIQUID CHROMATOGRAPHY DATA	214
APPENDIX D: DATA FOR ARTIFICIAL NEURAL NETWORKS.....	229

LIST OF FIGURES

Figure 1.1 Thesis Outline.....	9
Figure 2.3: Growth phases in batch fermentation (Shuler and Kargi, 2005)	26
Figure 2.4 Biomass versus time in a fed-batch process (Saarela <i>et al.</i> , 2003).....	29
Figure 2.5: Kinetics of ethanol production during fermentation (Vu & Le, 2010).....	30
Figure 2.6: Illustration of natural neuron in human.	43
Figure 2.7: Single Node Multi-Layer Perceptron (Patan, 2008)	46
Figure 2.8: Neural Network topology with a single hidden layer (Venkata <i>et al.</i> , 2009). ..	47
Figure 2.9: Energy consumption curve and possible energy recovery curve in a brewing process (Muster-Slawitsch, 2011).....	56
Figure 2.10: Flow sheet for the simple batch process (Halim <i>et al.</i> , 2009)	64
Figure 2.11: TAM and TSM representation (Halim <i>et al.</i> , 2009).....	67
Figure 3.1: Schematic Experimental set-up	73
Figure 3.2: Wits mini brewery plant (Madigoe, 2009)	73
Figure 3.3: (a) BIOSTAT fermentation system used for validation of ANN model, (b) Process flow diagram of the BIOSTAT System	79
Figure 3.4: HPLC Chromatogram of five Carbohydrate standards.	81
Figure 3.5: High Performance Liquid Chromatography used for sample analysis	82
Figure 3. 6: The Alcohol Gay-Lussac %(v/v) of standard mixtures of alcohol.	83
Figure 4.1: Measure of the proportion of the total variation between observed and predicted values.	95
Figure 4.2 Response Surface and contour plots at constant concentration of fermentable sugars	97
Figure 4.3 Response Surface and contour plots at constant yeast concentration	99
Figure 4.4 Response Surface and contour plots at constant temperature.....	101
Figure 5.1 Framework of the prototype for intelligent optimisation technique using neural networks and genetic algorithms.....	110
Figure 5.2 Train, Validation and Test data mean squared error against epoch.....	117
Figure 5.3 Plot of Target and Predicted values when using 2 neurons in hidden layer ..	119
Figure 5.4 Plot of Target and Predicted values when using 10 neurons in hidden layer	120
Figure 5.5 Plot of target and predicted values using 30 neurons in hidden layer	120
Figure 5.6 Architecture of the developed Neural Network Model	121
Figure 5.7 Validation plot for ANN model.....	122
Figure 5.8 Test data for ANN architecture	123
Figure 5.9 Genetic Algorithm population plot.....	125
Figure 5.10 Levels reached by the stopping criteria	128
Figure 5.11 Histogram of the parents.....	130
Figure 5.12 Average distance between individuals at each generation.....	131
Figure 5.13 Expected number of children vs. raw scores at each generation	132
Figure 5.14 Best, Worst and Mean fitness function value (i.e. Ethanol concentration) in each generation	133
Figure 5.15 Histogram of the scores at each generation	134
Figure 5.16 Scores of individuals at each iteration	135

Figure 6.1 Ethanol productivity for batch and fed-batch fermentation at 10°C	140
Figure 6.2 Ethanol productivity for batch and fed-batch fermentation at 20°C	141
Figure 6.3 Ethanol productivity for batch and fed-batch fermentation at 30°C	143
Figure 6.4 Ethanol productivity for batch and fed-batch at 35°C	144
Figure 6.5 Time event charts for the streams	152
Figure 6. 6Temperature profiles for batch process	154
Figure 6.7 Time event chart for rescheduled process	154

LIST OF TABLES

Table 2-1 Advantages and disadvantages of batch processes (Hewitt, 2007)	27
Table 2-2 Advantages and disadvantages of fed-batch processes (Hewitt, 2007)	30
Table 2-3 Experimental layout and runs for the central composite rotatable design	38
Table 2-4 Stream information for simple batch process in figure 2.....	65
Table 4-1 Relationship between coded and actual values of the variables (Napier-Munn, 2000)	88
Table 4-2 Experimental layout and runs for the central composite rotatable design	90
Table 4-3 Experimental Ethanol Yield	92
Table 4-4 ANOVA for the fitted model.....	94
Table 4-5 Regression coefficients for refitted model.....	94
Table 6-1 Standard deviations between experimental and theoretical data for batch and fed-batch fermentations	146
Table 6-2 Plant operating schedule.....	150
Table 6-3 Stream data for the case study	151

LIST OF SYMBOLS

Symbol/ Unit	Definition	Units
c	Product concentration	kg/m^3
D	Dilution rate	$1/\text{h}$
k	Proportionality constant	$1/\text{h}^2$
$m(\text{subscript})$	Maximum production rate	
r	Production rate	$\text{kg/m}^3\text{h}$
RDF	Real Degree of fermentation	%
O	Wort original gravity	SG
E	Real extract gravity	SG
y	Predicted response	
β_0	Coefficient of intercept	
β_i	Coefficient of linear effect	
β_{ii}	Coefficient of quadratic effect	
β_{ij}	Coefficient of interaction effect	
ε	Term that represents other sources of variability not accounted for by the response function	
X_i and X_j	Coded independent variables	
	Colony	units/ml
	number of colonies	
	dilution factor	
Q	total heat transferred	Joules
m_L	mass of liquid in the vessel	kg
C_{pL}	specific heat capacity of liquid	J/kg.K
T_{L2}	final temperature of the liquid in the vessel	K
T_{L1}	initial temperature of the liquid in the vessel	K

CHAPTER ONE

Introduction

“It's impossible to map out a route to your destination if you don't know where you're starting from”

- **Suze Osman**

1. INTRODUCTION

1.1 Introduction

Beer brewing is one of the most widely practiced processes in the world. The brewing technology is associated with huge business opportunities because of high demand for the final product (*i.e.* beer and other non-alcoholic beverages). For this reason, it is necessary to operate under conditions which can meet high demand for the product, and simultaneously reduce capital and operating costs. Currently most brewing companies use batch fermentation process and free floating yeast cells. These might be attributed to the fact that early processes were designed in that format and people are comfortable and reluctant to change. Also, batch mode is very easy to set up and operate; allows flexibility to produce multiple related products in the same facility; and is capable of handling variations in feed stock, product specifications and market demand pattern (Cinar, 2003). Besides these advantages, due to worldwide competition and increasing costs of resources, businesses have to consider innovative ways of brewing beer optimally.

Considerable interest has developed in the use of fed-batch fermentation as an alternative to batch fermentation in brewing because of the advantages that fed-batch fermentation presents (Vu and Lee, 2007a, 2010). For example, by controlling the substrate concentration in the reactor fed-batch allows the brewer to have control on the fermentation reaction and this helps to overcome several obstacles such as temperature control, oxygen transfer, *etc.* (Yi-Huang *et al.*,

2012). Substrate control also allows metabolic control, to avoid osmotic effects, catabolite repression, and overflow metabolism of side products. Fed-batch fermentation has been used successfully for many types of fermentation processes, such as lactic acid production (Liu *et al.*, 2005), biomass production (Demirci *et al.*, 1999), antibiotic production (Colombie *et al.*, 2005), and recombinant proteins production (Choi and Park, 2006). Fed-batch involves a process whereby fermentation material can be added or removed during the run. Although fed-batch fermentation has been successfully implemented in other industries, it has attracted low interest in the brewing industry. From the literature, positive fed-batch results are obtained with supplementation of feeding-medium with nutrients, which is not a welcomed practice by the beer consumers.

Latest trends in the chemical engineering industry and academia include the desire for a cleaner production processes coupled with an increased yield of the desired product. The past decades has seen significant efforts devoted to the development of methodologies that aim to reduce emissions and waste, reduce operating capital cost for a variety of chemical process industries (*e.g.* biochemical, petroleum, pharmaceutical, food and beverage industries, *etc.*). Although process optimisation and integration of brewing processes is attracting more and more attention due to the rising capital and utility costs, it is still a challenging task, particularly because the process exhibits strong non-linear behaviour (Youssef *et al.*, 2000). Additionally, their non-steady and non-ideal characteristics need to be taken into consideration. Process optimisation and integration is of great importance for research and practice in the brewing industry.

In recent years, studies based on the optimisation of ethanol production during fermentation have been undertaken. Fermentation is an important step in brewing so optimisation of fermentation will lead to an increase in the alcohol yield at a lowest possible cost. Biochemical systems, such as brewing have some components that are well defined, along with numerous others that are only partially understood (Triadaphillou *et al.*, 2005). Straightforward mathematical or algorithmic modelling in biochemical systems is only possible for some subsystems. As a result, new techniques need to be explored that can deal with biochemical process's complexities. The fermentation step involves a lot of variables, so in the process of optimising fermentation, one should use a method that will try and take into account the effect of all the variables. One widely used method in optimisation is the Response Surface Methodology (RSM) based on the Central Composite Design (CCD). RSM is widely used because it defines the effect of the independent variables alone or in combination on the process as stated by Bas and Bayoci (2007). With available experimental data one can generate a mathematical model that describes the chemical process that is being optimised. The research therefore, aim to investigate the effectiveness of Response Surface Methodology through formulation and optimisation of empirical model by using RSM in relating yeast concentration, temperature and media composition to ethanol production during fermentation.

Optimisation of fermentation processes is a challenging task that relies on an understanding of the complex effects of processing inputs on productivity and quality outputs (Adeyemo and Eniton, 2011). Due to the complexity of biological activities associated with fermentation, studies have shown that optimisation

methods based on utilising mathematical models and statistically designed experiments can be less effective in modelling fermentation processes since they are ineffective in predicting non-linear dynamic systems exhibited by biological processes and thus, there is a great need for new tools that can be employed for the optimisation of fermentation (Vlassides *et al.*, 2000). The importance of this study is reflected on the novel use of Artificial Neural Networks as an optimisation tool for biological systems because of its great potential to circumvent the limitations associated with the traditional optimisation methods used in the fermentation process and thus, it can serve as a basis of discernment in improving the brewing process.

Artificial neural networks are a powerful data modelling tool designed in such a way as to mimic the human learning processes by creating linkages between the process input and output data (Syu and Tsao, 2010). The power of neural networks lies in their ability to represent both linear and non-linear relationships and in addition, they learn these relationships directly from the data being modelled. Furthermore, they rely on inductive inference or generalisation to recognize invariance and capture higher order correlations and generate relationships from inputs and outputs from a set of examples (Becker *et al.*, 2002).

Inefficiency in energy optimisation of batch processes relative to the continuous processes as well as the increase in energy cost makes it necessary to save energy through other strategies such as process integration (Majozi, 2006). Process integration is a family of methodologies for combining several processes and/or

process units to reduce consumption of resources or harmful emissions to the environment. Process integration for energy saving and waste reduction is a subject of increasing interest among plants operating under batch mode. In brewing processes, which operates under batch mode, cold and hot stream are scarcely simultaneous. This makes it difficult to integrate processes through direct heat exchange (Bozan *et al.*, 2001, Puigjaner *et al.*, 2007). If process integration can be applied inter or intra processes, it has the potential to make the operation more profitable and economical. Hence, the research has also looked into benefits of batch process integration in the brewing industry.

1.2 Research problem and contribution to knowledge

Although various optimisation techniques have been applied in other chemical processes, their application in biochemical industry, more specifically brewing, has been limited as there is little data available. Limited application of process integration and optimisation method to the brewing industry might be attributed to the fact that biochemical processes have complex systems whose behaviour is difficult to model (Batunde, 2007). Also, their time-based mode of operation (batch) makes it difficult to integrate. However, many authors have agreed that, advanced optimisation techniques such as Response Surface Methodology (RSM) and Artificial Neural Networks (ANN) can be successfully deployed to optimize complex processes. It is also theorised, that, time-based integration methods can be used for batch processes integration. These are the fundamental concepts that the research is based on.

As pointed out previously, slow response nature and non-linear behaviour of brewing process makes it difficult to optimise using traditional optimisation techniques. The research aims to provide a new and efficient method for optimizing brewing processes using Artificial Neural Networks (ANN) in comparison to Response Surface Methodology (RSM). The research aims to also provide efficient integration method for batch brewing fermentation process. New directions into the research of fed-batch fermentation should also be introduced.

1.3 Aim and objectives

The main aim of this study is to develop ANN model for optimisation of brewing process, compare ethanol production yield between batch and fed-batch fermentation in order to evaluate the prospects of switching to fed-batch fermentation and integrate the process to optimize resources. The objective functions were ethanol yield and production rate.

The specific objectives are:-

- i. Develop an Artificial Neural Network (ANN) model that relate selected process variables to performance for complex brewing fermentation systems;
- ii. Investigate the reliability of RSM for analysis and optimisation of selected variables and compare it to ANN;
- iii. Comparison of fed-batch fermentation and batch fermentation; and
- iv. Integrate process units and production processes the Wits micro-brewery.

1. 4 Research Method

The research method of this study involved: - literature review, experimental design, laboratory work, data analysis, conclusions, recommendations and documentation.

1.5 Thesis Outline

The thesis is spread out into seven chapters; that are schematically represented in Figure 1.1.

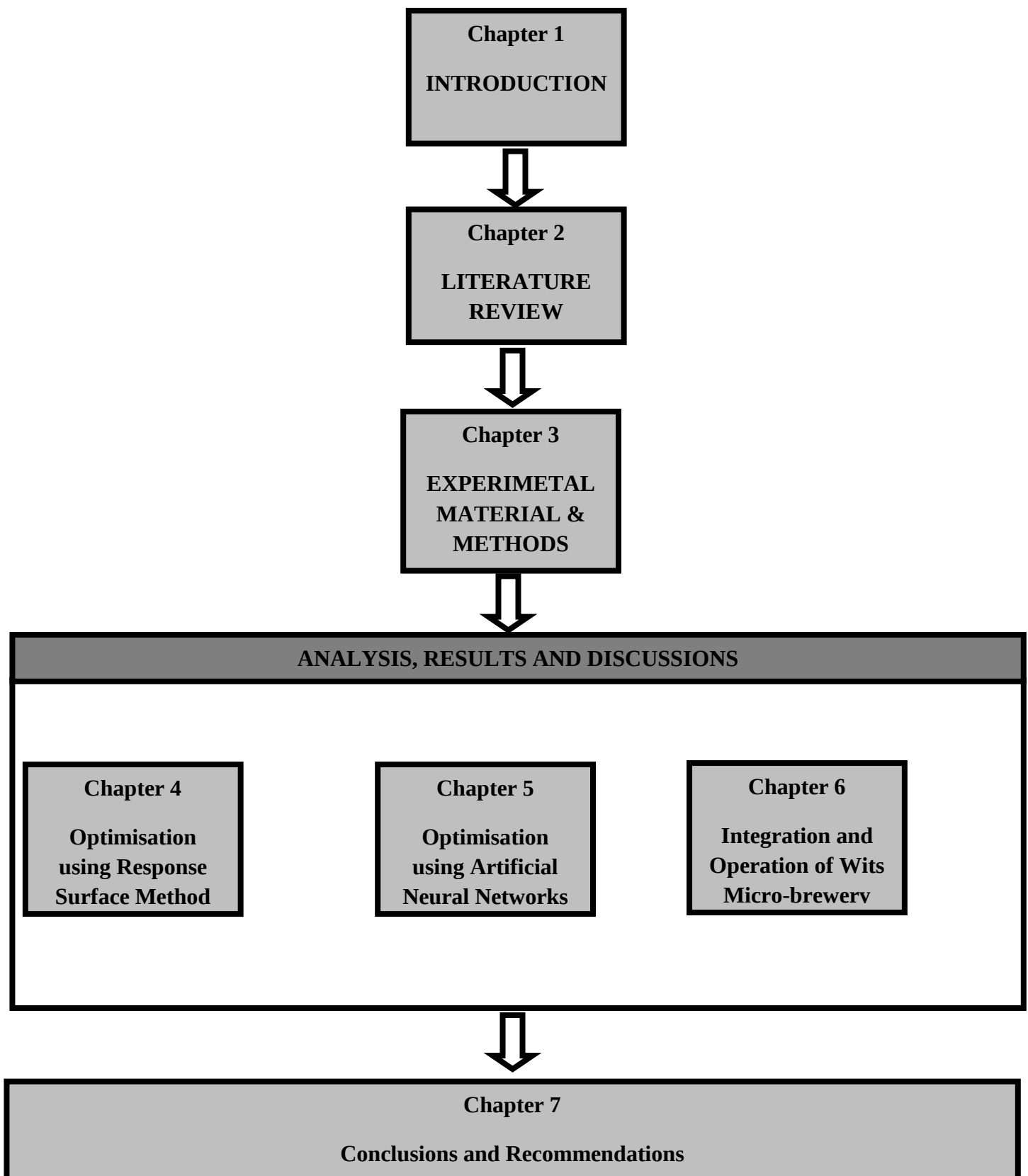


Figure 1.1 Thesis Outline

Chapter 1

Chapter one provides motivation and background for the research. The problem statement, overall objectives of the study and brief description of research methodology are also discussed in this chapter.

Chapter 2

Brewing industry is reviewed to provide insight into the latest progress in the development of the techniques involved. The evolution of the brewing industry is covered to give the background of the industry. General knowledge of the brewing process and the general principles are also discussed. Different types of fermentation processes are reviewed, together with models that govern them. Challenges in optimisation of brewing process are highlighted and current status is given. Key points in optimisation methods such as Response Surface Methodology (RSM), Artificial Neural Networks (ANN) and Genetic Algorithms (GA), which are relevant to the focus of this research, are revised with additional observations during the course of the research. Emphasis is placed on the superiority of ANN as compared to other traditional methods.

Challenges in integration of batch processes are highlighted also. Integration methods that have been theorised and/or used by other researchers and their capability in integrating batch processes are demonstrated.

Chapter 3

In this chapter, materials and methods that are used in this study are described. Data analysis and validation methods employed in the study are demonstrated. Statistical analysis used to ensure reliability of the results is also demonstrated.

Chapter 4

The chapter demonstrates the capability of RSM to solve complex systems such as brewing process. Because of the difficulty in modelling ANN, especially for complex systems, the objective of this chapter was to investigate the performance of optimising brewing process using relatively less complex methods such as RSM. Laboratory tests, findings and conclusions from the study are reported. The capability and limitations in using RSM are demonstrated.

Chapter 5

The chapter proposes a new reliable hybrid optimisation method for brewing process that combine ANN and Genetic Algorithm. Superiority of the ANN over other traditional optimisation methods including RSM is demonstrated in this chapter. The reliability of the new method is supported by experimental tests.

Chapter 6

This chapter discusses integration and operation of the Wits micro-brewery. Basic comparison of batch and fed-batch fermentation is carried out in this chapter. For the first time, this comparison is carried out without any fermentation stimulus supplements, such as gluconates (alpha/beta) and phosphates as it has been the case with other researchers. The objective was to investigate the possibility of switching to fed-batch fermentation in the Wits Micro-brewery plant. Industrial limitations of fed-batch fermentations are demonstrated. Besides from the optimisation of brewing process units such as lauter tun, mashing, boiling, and fermentation, additional benefits can be realised if the process units are integrated for energy saving. Capability of process integration method, rescheduling, is demonstrated in this chapter. Potential benefits that can accrue from process integration are reported.

Chapter 7

This chapter concludes the thesis by summarising the conclusions of all the findings in the study. Relevance of the thesis is confirmed by verifying that all objectives have been met. Possible areas for future researches that can possibly follow this study are recommended. Potential application of the findings in the industry is also demonstrated. Appendices of other relevant information are also included.

CHAPTER TWO

Literature Review

“Information is not knowledge”

- **Albert Einstein**

2. LITERATURE REVIEW

2.1 General Introduction

Beer brewing industry is a large and sophisticated industry that incorporates agricultural and industrial sectors. As an employer to thousands of people, it plays a major role in the economy. In recent years, industries across the globe have seen a sharp increase in input and operational costs, this was mostly notably in energy costs and feed inputs. World population, and hence demand for the product is significantly increasing, with most demand reported in developing countries and Asian market (Andrews *et al.*, 2011). Due to environmental concerns and depleting resources, scientists have a challenge to come up with optimised and cleaner ways of producing beer with minimal resources and less harm to the environment, while keeping up with the demand.

In this literature review, the brewing industry is reviewed to provide insight into the latest progress in the development of the techniques involved. The evolution of the brewing industry is covered to give the background of the industry. General knowledge of the brewing process and the general principles are discussed. Different types of fermentation processes are reviewed, together with models that govern them. Challenges in optimisation of brewing process are highlighted and the current status is given. Key points in optimisation methods such as Response Surface Methodology (RSM), Artificial Neural Networks (ANN) and Genetic Algorithms (GA), which are relevant to the focus of this research, are revised with

additional observations during the course of the research. Emphasis is placed on the superiority of ANN as compared to other traditional optimisation methods.

Challenges in integration of batch processes are highlighted. Integration methods that have been theorised and/or used by other researchers and their capability in integrating batch processes are demonstrated.

2.1.1 Brewing Industry

Beer is generally considered to have originated as a by-product of the development of agriculture, although a minority opinion holds that the cultivation of cereals originated as a consequence of man's desire for alcohol rather than vice versa. Surviving historical artefacts allow us to trace brewing back to the Mesopotamians around 6000 or 7000 years ago (Priest and Stewart, 2006, John, 2005). The Ancient Egyptians where brewers brewed from the indigenous cereal sorghum is still integral to the politics of African tribal life. In northern Europe cold conditions inhibited the development of viticulture; it is historically found that there is a link between development of agriculture and the brewing industry (Wilson and Gourvish, 1993). In the tenth century, the use of hops in brewing spread from Germany across Europe to replace or at least supplement the plethora of plants, herbs and spices popular at that time (Priest and Stewart, 2006).

Historically, the terms ale and beer were being applied to distinct beverages made by separate communities of brewers. Ale described the drink made without hops and the term beer was reserved for the hopped beverage. By the 16th century, ale

brewers came to use hops in their brews, but at a lower level than was usual for beer and thus a distinction remained between beer and ale. Ale would be recognised as a heavy, sweet, noticeable alcoholic drink characteristic of rural areas, whereas beer was bitter, often lighter in flavour and less alcoholic, but frequently darker brown in colour than ale and was popular in towns (Hornsey, 1999a, 2004).

In the 19th century, industrialisation, population growth and urbanisation resulted in increased consumption of beer (Adams *et al.*, 1995). In the leading European beer-drinking countries: the United Kingdom, Germany, and Belgium, there was a two-to four-fold increase in output between 1830 and 1900 (Wilson and Gourvish, 1993). The development of railway networks from the 1830s transformed distribution, prompting the larger scale producers to make their regional beers available internationally to a growing population.

By the 20th century, the top three brewing countries, Germany, the United Kingdom and the United States accounted for 68.5% of recorded beer production (Anderson, 2003). During the first half of the 20th century, although individual breweries and maltings had grown larger and outputs had increased, there was little increase in batch or vessel size. The use of stainless steel, greater chemical, physical, and biochemical knowledge, better analytical control, increased availability of process aid and then automated computer-controlled plants, led to step changes. The brewing industry is currently a huge global business.

2.2 Beer Brewing Process

2.2.1 Background

There are several steps involved in the beer brewing process, which include, malting, milling, mashing, lautering, boiling, fermenting, conditioning, filtering, and packaging. Although the ingredients used in beer brewing might have barely changed from those used centuries ago, process units designs and operational strategies have evolved over time. Current research in brewing technology is focusing on optimisation of process units in order to gain maximum return from each batch stage. Breweries are looking into implementing technologies that have successfully been applied to other biotechnology processes such as pharmaceuticals, waste water treatment and other food processing industries.

2.2.2 Malting

Barley is the preferred primary raw material for brewing process. Barley grain is unacceptable for direct use in the brewing process in its native form, the modification that is generally done to make it suitable for brewing is known as malting (Reed and Nagodawithana, 1991). In this process raw barley is converted by controlled steeping, germination, and kilning into a product that is much more friable, with increased enzyme levels and with altered chemical and physical properties as stated by Goldhammer (2008). In this process, the carbohydrates available for the development of the embryo undergo minimal modification by enzymes developed within the seed.

The grain is conveyed into steep tanks, which are filled with water and are equipped with aeration device to aerate the wetted grain (Briggs *et al.*, 1981a; Briggs, 1998). The steeping process is done until the moisture content of the grain reaches 42-45%. The grain is spread on germination rootlets which allow rootlets to form, the starches within the grain break down into shorter lengths and at this point the grain is called green malt (Briggs, 1998, Madigoe, 2009). The next stage in the process is kilning. In the kilning process, the germinated or green barley is dried to low moisture content under controlled conditions to preserve the enzymatic activity developed during the germination. The kilning chambers are similar to the germinating chambers except that heated dry air instead of cool humid air is passed through the grain bed, the rootlets formed during germination become dry and brittle during kilning and are removed on a scalper (Linko *et al.*, 1998). The finished product is now called malted barley; the malted barley is finally cleaned and conveyed to the storage bins and ready for milling.

2.2.3 Milling

Malted barley is milled to smaller particles to expose the inner starch. The milled malted barley is now called grist (Hough, 1985). In milling the malted barley, it makes the malt easier to wet at mashing stage and aids faster extraction of the soluble components from the malt during the enzymic conversion (Boivin *et al.*, 1997). Once the barley is malted, it can now be taken to the first stage in the brew house that is called mashing.

2.2.4 Mashing

Mashing is the process of mixing the grist with hot water to form a mash and letting the mixture stand for about one to three hours while the enzymes degrade the proteins and starch to yield a soluble malt extract, called the wort (Michael-Eskin, 1990). This stage also includes the separation of the wort from the solids.

2.2.5 Lautering

The wort (sweet liquid) is extracted from the mash in the lauter tun; the sweet wort is rinsed from within the wort by spraying the mash with hot water from the top (Briggs, 1998, Madigoe, 2009). The main aim of the lauter tun is to produce a high yield of sweet extract and less of the undesired malt constituents. Industrially, mashing and lautering happen in separate tanks. But at Wits Micro-brewery plant, both mashing and lautering takes place in a single tank at consecutive times as illustrated in Figure 2.1. At the end of the lautering process, the spent grains are discharged, and the wort is fed into boiler tank for further processing.

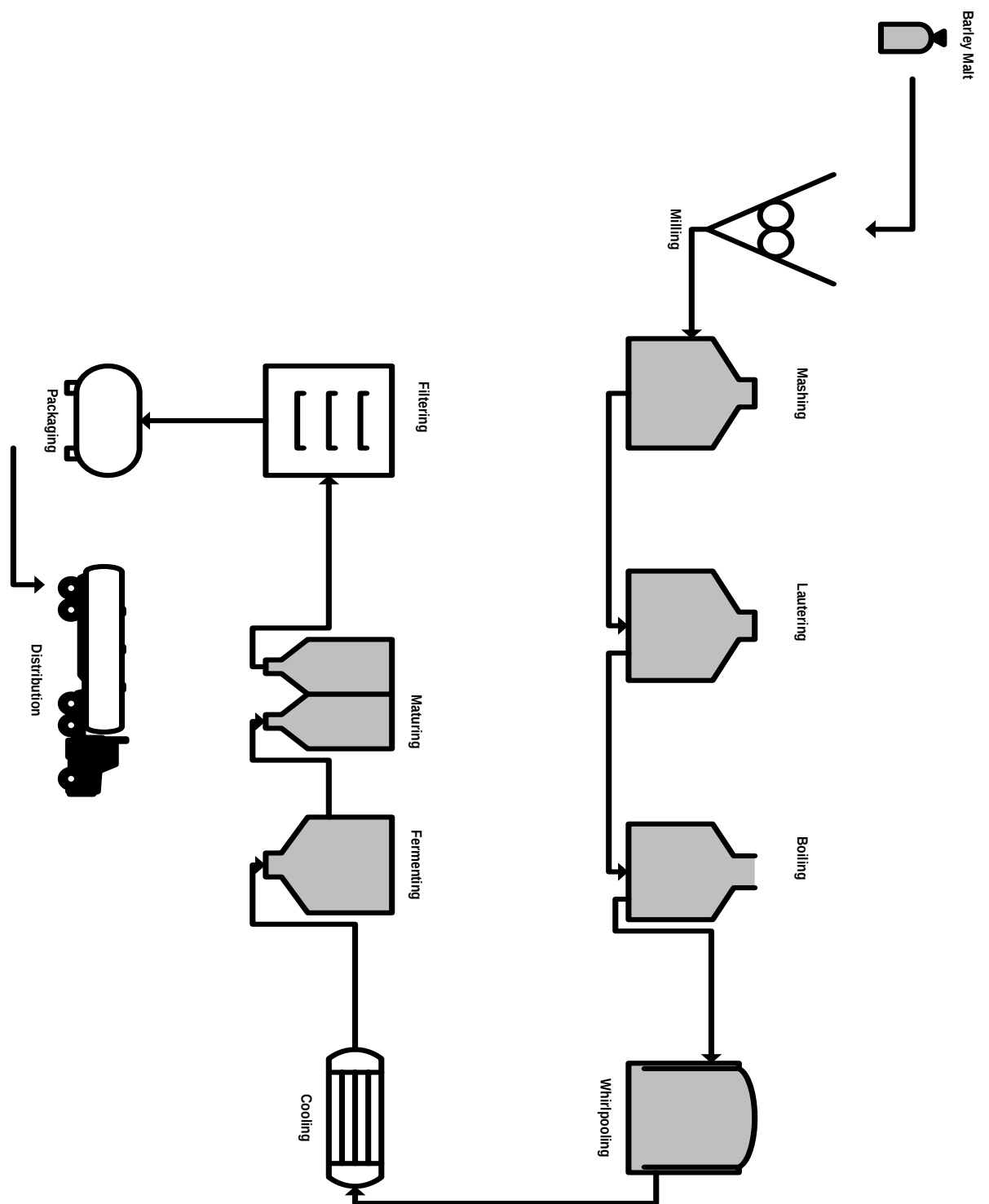


Figure 2.1: Typical brewing process

2.2.6 Boiling

Hough (1985) states that at this stage, the wort is boiled with hops to stop enzyme action, sterilises the wort and also to coagulate some proteins. The hops are for imparting distinctive flavour and aromas to the wort. After boiling, the wort is then transferred into the whirlpool. Whirlpooling is the process of settling coagulated matter, including proteins to settle at the centre of the tank, thereby purifying the wort even further. The wort is then clarified, cooled and aerated so that it is an ideal medium for yeast growth and fermentation. The wort is cooled to fermentation through heat shock process using cross flow heat exchanger and cooling water system as shown in Figure 2.1.

2.2.7 Fermentation

The clear wort that was quickly cooled through heat exchanger to fermentation temperature is now taken into the fermentation tank. The fermentable sugars in the wort are converted to ethanol and carbon dioxide, using a specially selected strain of brewing yeast (Harrison, 2009). The yeast uses up available oxygen and wort nutrients to increase in numbers by cell division. Top yeasts ferment at temperatures in the range 15-20°C and with bottom yeasts ferment at temperatures in the range of 7-14°C (Priest and Stewart, 2006, Briggs *et al.*, 1981a, Briggs, 1998). When fermentation is complete, the beer can optionally go through maturing and filtering. During maturing stage, the beer is conditioned around freezing point. Beer becomes smoother because unpleasant flavors cannot

dissolve in near-freezing point beer. Filtering the beer helps with stabilization of the flavor (Briggs, 1998).

2.3 Challenges in optimisation of brewing fermentation

2.3.1 Background

In the brewing context, fermentation is the production of alcohol from sugars through the action of microbes or by the action of enzymes. It is the heart of the brewing process. Fermentation is prone to alterations stemming from the variation in the leaving yeast organism and such might cause off-flavours in beer (Batunde, 2007). Yeast cells are used in fermentation to consume sugars and release ethanol. But yeast cells lose this converting ability as the fermentation process continues and as they are replenished, until they finally denature. Different yeast cells denature and lose this “ethanol production” ability at different rates, since many yeast cells are used during fermentation, even on a micro scale, this makes the entire fermentation process difficult to model (Triadaphillou *et al.*, 2005, Youssef *et al.*, 2000, Niu *et al.*, 2011). Also, the variations that come with using living and complex organisms may affect the fermentation speed or the final fermentation product taste (Pollanen *et al.*, 2001). Both the fermentation speed and the flavour of the beer are important factors in meeting the product specifications and might hence affect sales and so on. Hence there is a need for computational tools (like ANN on MATLAB) that will efficiently predict the course of the fermentation

process. The results from this research project hope to explore this rather robust area, making it applicable for industrial processes and profitability.

Several types of techniques have traditionally been used for the optimisation of fermentation processes and these are methods based on single variable optimisation, mathematical modelling, and response surface methodology (Vlassides *et al.*, 2000). In single variable optimisation, one process parameter at a time can be varied and an objective function is measured to get the value of the parameter that gives the best result. However, analysis of one parameter at a time when multiple parameters are to be optimised can lead to inefficient use of resources and it seldom guarantees the determination of optimal conditions. Moreover, it is laborious and time consuming (Dasari *et al.*, 2009). On the other hand, the main shortcoming of mathematical modelling lies in its inability to conceive non-linearity behaviour exhibited by biological activities during fermentation. With ANNs, several factors (parameters) can be simultaneously varied and they have emerged as an attractive tool for non-linear multivariate modelling (Comrie, 1997). The potential supremacy of ANN is that it is generic in structure and possesses the ability to learn from historical data, as already explained above. In comparison to the widely used Response Surface Methodology, ANNs does not require a prior specification of suitable fitting function and furthermore, it has a universal approximation capability, that is, it can model almost all kinds of non-linear functions whereas RSM is useful only for quadratic approximations (Kiran *et al.*, 2008; Dasari *et al.*, 2009).

2.3.2 Different types of fermentation processes

Challenges have been encountered due to the complexities associated with the fermentation process, a biochemical process with living organisms that behave sporadically. In simple terms, fermentation is the stage during brewing in which the wort is brought into contact with the yeast cells. The yeast cells ferment the sugars to produce ethanol, carbon dioxide and other less important by products. This process is governed by the following chemical equation:



The *Saccharomyces Cerevisiae* is the most commonly used yeast in brewing because it influences the quality and quantity of the secondary products formed during the fermentation process (Manzano *et al.*, 2005). In addition to that it has an advantage of being able to ferment sugars that other yeasts cannot ferment, specifically dextrins (Boswell *et al.*, 2003). Other ethanol production processes make use of aerobic fungi such as the common moulds of the general *Aspergillus*, *Fusarium* and *Mucor*.

There are different types of fermentation processes in brewing. Fermentation processes are classified according to the type of yeast cells used. The two mostly used types of fermentations are ale and lager. Ale uses the top cropping *Saccharomyces Cerevisiae* yeast at temperature between 14°C and 25°C, lager uses the bottom cropping *Saccharomyces Uvarum* between 8°C and 13°C (Cinar,

1999). The yeast uses the energy obtained during fermentation to multiply out and this is perfectly done in a neutral or slightly acidic environment.

There are a number parameters affecting the fermentation process such as temperature, metallic ions, oxygen content, pH, yeast concentration and concentration wort (sugars) (Bamforth, 2003). Metallic ions such as Mg^{2+} , Zn^{2+} , K^+ and Co^{2+} are vital to the lives of the yeast cells as it serves as nutrients, most of these cationic elements are present in the wort (Hornsey, 1999, 2004). Out of afore mentioned parameters, temperature is the most important one because it affects the fermentation rate as well as the beer flavour (Dragone *et al.*, 2008). The fermentation of beer occurs at low temperatures; generally on the range of 5°C to 25°C and this temperature range incorporates the fermentation of lagers and ales (Baxter and Hughes, 2001).

2.3.3 Batch and Fed-Batch Fermentation

The term "batch process" is often used to refer generically to both batch and fed-batch operations in fermentation. In the former case, the required ingredients used for fermentation are fed to the processing vessel at the beginning of the operation and no addition or withdrawal of material takes place during the batch run, whereas in the fed-batch fermentation process material can be added or removed during the run.

Fed-batch strategy is typically used in bio-industrial processes to reach a high cell density in the bioreactor. Mostly the feed solution is highly concentrated to avoid dilution of the bioreactor (Lee *et al.*, 1999). In most cases the growth-limiting nutrient is glucose which is fed to the culture as highly concentrated glucose syrup. Since both overfeeding and underfeeding of nutrient is detrimental to cell growth and product formation, development of a suitable feeding strategy is critical in fed-batch cultivation.

In batch bioreactor all components, except gaseous substrates such as oxygen, pH-controlling substances and antifoaming agents, are placed in the reactor in the beginning of the fermentation. During the process there is no input nor does output flow. When microorganisms are growing in a batch reactor certain phases of growth can be detected, as shown in a typical growth characteristic curve of Figure 2.2. It is worth noting that the appearance and the length of each phase depend on the type of organisms and the environmental conditions.

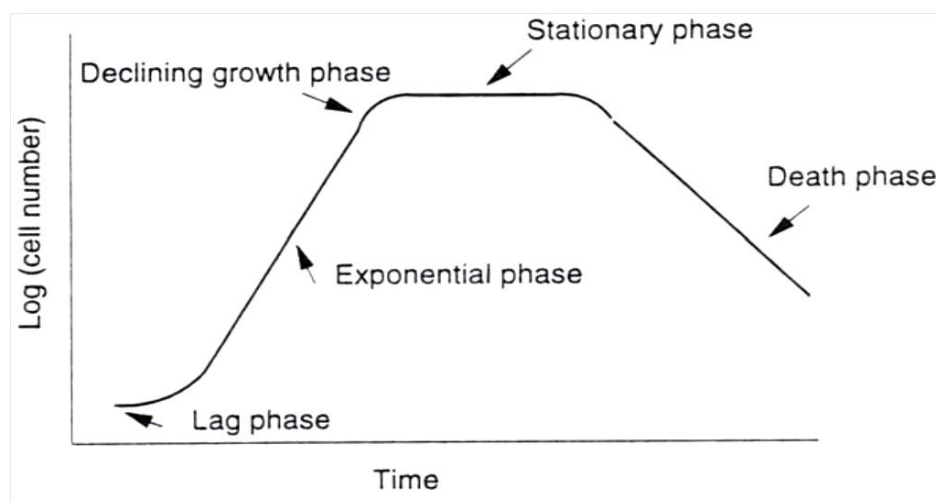


Figure 2.3: Growth phases in batch fermentation (Shuler and Kargi, 2005)

The first phase in the growth, where the growth rate stays almost constant, is the lag phase. The lag phase is caused for many reasons. For example, when the cells are placed in fresh medium, they might have to adapt to it or adjust to the medium before they can begin to use it for growth (Saarela *et al.*, 2003). When a substrate begins to limit the growth rate the phase of the declining growth begins. The growth rate slows down until it reaches zero and the stationary phase begins. In the stationary phase the number of the cells remains practically constant, but the phase is important because many products are only produced during it. The last phase is called the death phase. During the death phase the cells begin to lyse and the growth rate decreases.

Table 2-1 outlines a number of advantages and disadvantages of batch fermentation as outlined by Hewitt (2007).

Table 2-1 Advantages and disadvantages of batch processes (Hewitt, 2007)

Advantages	Disadvantages
• Reduced risk of contamination or cell mutation, due to a relatively brief growth period; • Lower capital investment when compared to continuous processes for the same bioreactor volume;	• Lower productivity levels due to time for filling, heating, sterilizing, cooling, emptying and cleaning the reactor; • Higher costs for labour and/or process control for this non-stationary process;

-
- | | |
|--|---|
| <ul style="list-style-type: none"> • More flexibility with varying product/biological systems; and • Higher raw material conversion levels, resulting from a controlled growth period. | <ul style="list-style-type: none"> • Increased focus on instrumentation due to frequent sterilization; and • Larger industrial hygiene risks due to potential contact with pathogenic microorganisms or toxins. |
|--|---|
-

Fed-batch fermentation is distinguished from batch fermentation by the addition, during operations, of a certain amount of fresh substrate and by the consequent withdrawal of a proportioned amount of broth (Longobardi, 1994). The fed-batch strategy is typically used in bio-industrial processes to reach a high cell density in the bioreactor (Neubauer and Ruottinen, 2008). Mostly the feed solution is highly concentrated to avoid dilution of the bioreactor. The controlled addition of the nutrient directly affects the growth rate of the culture and allows avoiding overflow metabolism (formation of side metabolites, such as acetate for *Escherichia coli*, lactic acid in cell cultures, and ethanol in *Saccharomyces Cerevisiae*) and oxygen limitation (anaerobiosis). In most cases the growth-limiting nutrient is glucose which is fed to the culture as highly concentrated glucose syrup (600-850 g/l) (Runquist *et al.*, 2009).

Figure 2.3 below represents biomass concentration as the function of time in a typical fed-batch process (Saarela *et al.*, 2003). The process is at first started as a batch process, but it is exhibited from reaching steady state by starting substrate feed once the initial glucose is consumed. The fermentation is continued at a certain growth rate until some practical limitation inhibits the cell growth. Different types of substrate limitations can be used in the fed-batch processes. The repression of the growth rate can be achieved for example by sugar, nitrogen or phosphate sources (Kumar *et al.*, 2000). If no reaction rate control is used, and the cells are growing exponentially, the reaction will eventually be limited by oxygen or by heat, as demonstrated in Figure 2.4 (Saarela *et al.*, 2003).

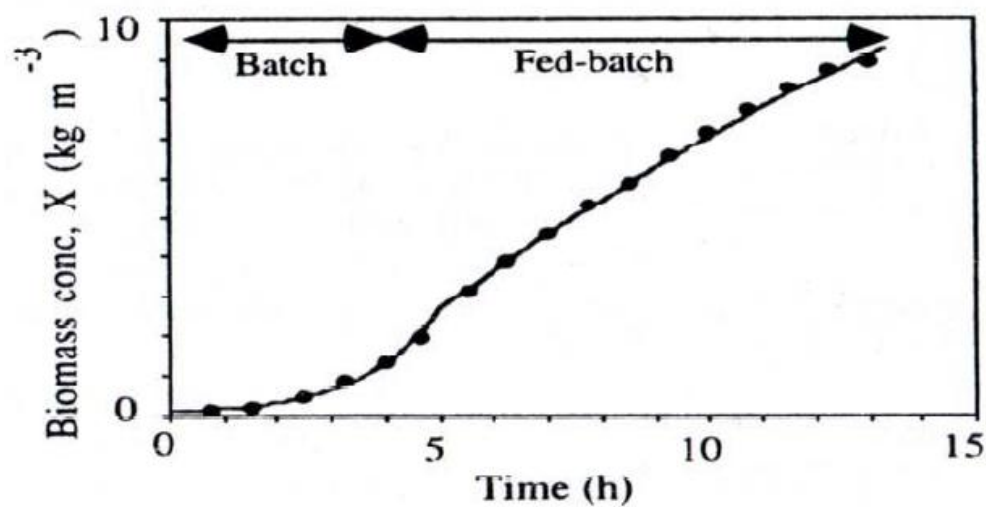


Figure 2.4 Biomass versus time in a fed-batch process (Saarela *et al.*, 2003)

In order to give the best possible growing conditions the pH and temperature levels are usually kept constant (Gregersen and Jorgensen, 1999), as fermentation systems are very sensitive to abnormal changes in operating conditions.

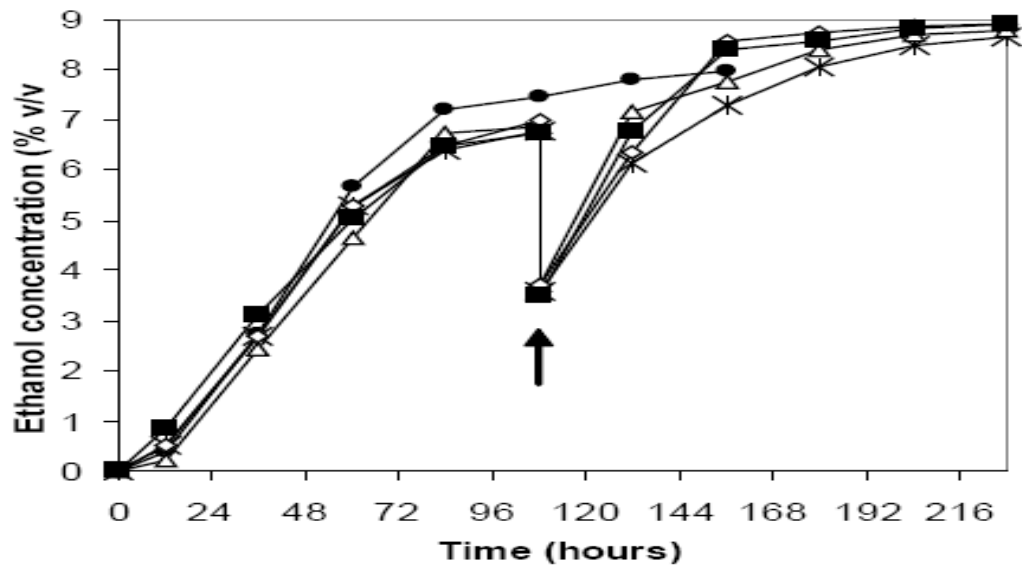


Figure 2.5: Kinetics of ethanol production during fermentation (Vu & Le, 2010).

In figure 2.4, the round-dotted line is batch fermentation whereas the rest represent different cultures in fed-batch fermentation.

Table 2-2 outlines a number of advantages and disadvantages of batch fermentation as outlined by Hewitt (2007).

Table 2-2 Advantages and disadvantages of fed-batch processes (Hewitt, 2007)

Advantages	Disadvantages
- Higher yield, resulting from a well-defined cultivation period during which no cells are added or removed, - Increased opportunity for optimizing environmental	Lower productivity levels due to time-consuming procedures for filling, heating, sterilizing, cooling, emptying and cleaning the reactor, and

conditions of the microorganisms
in regard to the phase of growth or
production and age of the culture,
and

- Greater expenses in labour and/or dynamic process control for the process.
 - Nearly stationary operation, important with slightly mutating microorganisms and those at risk for contamination.
-

2.3.4 Fermentation Modelling

Fed-batch fermentation is distinguished from batch fermentation by the addition, during operations, of a certain amount of fresh substrate. In order to have an idea for the substrate addition rate, it is of importance to know the rate of production of ethanol. The following equation is the theoretical estimation of ethanol production rate (r) for the batch mode of fermentation:

$$r = r_m \left(1 - \left(\frac{K(c - c_m)^2}{2r_m^2} \right) \right) \quad (2.2)$$

where r_m is the maximum production rate, K is a proportional constant specific for each fermentation process, c is the production concentration and c_m is the maximum production concentration (Longobardi, 1994). But the production rate is defined as follows:

$$r = \frac{dc}{dt} \quad (2.3)$$

Equation 2.3 can therefore be rewritten as follows:

$$\frac{dc}{dt} = r_m \left(1 - \frac{K(c - c_m)^2}{2r_m^2} \right) \quad (2.4)$$

For fed-batch the following equation is the theoretical estimation of the ethanol production rate:

$$\frac{dc}{dt} + Dc = r_m \left(1 - \frac{K(c - c_m)^2}{2r_m^2} \right) \quad (2.5)$$

It can be noted that the only difference between Equation 2.4 and Equation 2.5 is the Dc term; D is the dilution rate (h^{-1}). Dc is defined as the withdrawal flow rate divided by the broth volume inside the fermentation tank. Full derivation can be obtained from Longobardi (1994), and other authors working on the same material.

In order to measure the amount of sugar fermented one can use the Real Degree of Fermentation (RDF) equation given below (Priest and Stewart, 2006).

$$\frac{O - E}{O - 1.064} = \text{RDF} \quad (2.6)$$

where the O is the wort original gravity and E is the real extract gravity.

The main reactions that take place in the fermentation process are given below (Priest and Stewart, 2006):

(2.7)

(2.8)

2.4 Response Surface Methodology

The high degree of complexity in brewing technology and other fermentation technologies usually makes it difficult for systematic approach and impossible or difficult to apply simple optimisation techniques. Current production costs of beer are higher and hence require further developments to reduce production costs. This can be achieved through optimisation of process parameters and/ or process units. Response Surface Methodology (RSM) has been successfully deployed in other industries and has demonstrated potential results in the food technology as well.

2.4.1 Modelling Response Surface Methodology (RSM)

Response Surface Methodology (RSM) is a mathematical and statistical technique which is useful for the modelling and analysis of problems in which a response of interest is influenced by several significant variables and the objective is to optimise this response (Bas and Boyaci, 2007, Simate, 2009). RSM attempts to analyse the effect of the independent variables, alone or in combination, on a

specific dependent variable (response). Furthermore, RSM also generates a mathematical model. The graphical representation of the model has led to the coined term Response Surface Methodology (Myers and Montgomery, 2002). The independent variables are presumed to be measurable and continuous, and can be controlled with negligible error, whereas the response is postulated to be a random variable (Oraon *et al.*, 2006). The practical application of RSM requires developing an approximating model for the true response surface. The second-order response surface model that might describe the relationship is (Wu *et al.*, 2009, Myers and Montgomery, 2002).

(2.9)

In Equation 2.9 above, y is the predicted response, β_0 is the coefficient of the intercept, β_i is the coefficient of the linear effect, β_{ii} is the coefficient of the quadratic effect, β_{ij} is the coefficient of interaction, k is the number of variables, and x_i are the coded independent and ε is the statistical error. The quadratic model has a constant term, k -first order terms, k -quadratic terms, and $k(k-1)/2$ interaction terms and thus, has a total of $p = (k+1)(k+2)/2$ terms (Draper and Lin, 1990). The model given (Equation 2.9) is written in matrix notation as (Myers and Montgomery, 2002).

(2.10)

where

$$, \quad , \quad , \text{ and} \quad (2.11)$$

In general, Y is an $n \times 1$ vector of the observation model matrix consisting of the levels of the independent variables expended to model form, β is a $p \times 1$ vector of the regression coefficients, and ε is an $n \times 1$ vector of random errors.

The equations given above are solved using the method of least squares which is a multiple regression technique. The difference between the observed and the fitted values for the i^{th} observation is called the residual and is an estimate of the corresponding , where is the residual, is the observed value, and is the predicted value (Cowpe *et al.*, 2007).

The criterion for choosing the regression coefficient estimates () is that they should minimise the sum of the squares of the residuals, which is often called the sum of squares of the errors ().The function is (Bas and Boyaci, 2007).

$$(2.12)$$

The residuals may be written as (see Eq. 2.10).Therefore, the least squares estimated of the elements of in Equation 2.13 are given as (Marchitan *et al.*, 2010, Myers and Montgomery, 2002).

(2.13)

where X' is the transpose matrix of X and $(X'X)^{-1}$ is the inverse matrix of

After the regression coefficients have been obtained, the estimated response can then be easily calculated by substituting the coefficients into Equation 2.9.

2.4.2 Central Composite Rotatable Design

The experimental design techniques commonly used for the process analysis and modelling of response surfaces for second order designs are the full factorial, partial factorial and Central Composite Rotatable Designs (CCRD) (Obeng *et al.*, 2005). If the design is rotatable as in the case of the CCRD, the variance of the predicted response, y , remains constant at all points, which are equidistant from the design centre and is not a function of direction (Simate, 2009; Myers and Montgomery, 2002, Khuri and Cornell, 1987). CCRD is more useful than the full factorial since it requires much fewer tests, but is shown to be sufficient to describe the response (Obeng *et al.*, 2005, Aslan, 2007).

The CCRD involves the use of three portions, (i) 2^k factorial, portion with its origin at the centre (conventionally called a cube with ± 1 being the factorial levels) allowing estimates of first order effects or first order two factor interactions; (ii) $2k$ axial or star points fixed axially at a distance α from the origin representing a one factor-at-a-time portion designed to generate the pure quadratic

terms, λ , in Eq. 4 (Simate, 2009, Myers and Montgomery, 2002); (iii) and n_c replicate points at the centre to provide information about the existence of the curvature in the system and stabilizes, the variance of the predicted response and provide an independent estimate of experimental error (Myers and Montgomery, 2002, Obeng *et al.*, 2005). Thus, there is a total of 2^k (or $2^{(k-q)}$) $+2k + n_c$ points in the design; where k is the number of factors to be studied and the $-q$ represents the fraction of the number of factors (where $q = 0$ for full factorial design) (Myers and Montgomery, 2002, Deyhimi *et al.*, 2011).

The axial points are chosen so that they allow rotatability which ensures that the variance of the model prediction is constant at all points equidistant from the design centre (Box and Hunter, 1957). A value of λ guarantees rotatability despite the number of centre points where 2^{k-q} specifies a two level design on k factors. The number of centre point replications, n_c , can also be chosen to cause a CCRD to have the variance of the predicted response, y , remaining constant at all points which are equidistant from the design centre. For three factors, the centre points are calculated as follows (Khuri and Cornell, 1987).

$$- \quad (2.14)$$

The units of the independent variables differ from one another. Since parameters have different units and/or ranges in the experimental domain, the regression analysis should only be performed once the parameters have been normalized.

Once the desired ranges of values of the variables are defined, they are then coded. Coding ensures that all the variables affect the response more evenly, and so, the units of the parameters are irrelevant. The relationships between coded and actual values of the variables in a CCRD were logically set as given in Table 2-3.

Table 2-3 Experimental layout and runs for the central composite rotatable design

Standard Runs	Coded values of variables			Actual levels of variables		
	Temperature	Sugars	Yeast	T(°C)	Sugars(g/500ml)	Yeast(g/500ml)
Factorial Points						
1	-1	-1	-1	10	10	0.2
2	1	-1	-1	19	10	0.2
3	-1	1	-1	10	30	0.2
4	1	1	-1	19	30	0.2
5	-1	-1	1	10	10	0.4
6	1	-1	1	19	10	0.4
7	-1	1	1	10	30	0.4
8	1	1	1	19	30	0.4
Axial Points						
9	-1.682	0	0	7	20	0.3
10	1.682	0	0	22	20	0.3
11	0	-1.682	0	14.5	wort	0.3
12	0	1.682	0	14.5	40	0.3
13	0	0	-1.682	14.5	20	0.1
14	0	0	1.682	14.5	20	0.5
Centre Points						
15	0	0	0	14.5	20	0.3
16	0	0	0	14.5	20	0.3
17	0	0	0	14.5	20	0.3
18	0	0	0	14.5	20	0.3
19	0	0	0	14.5	20	0.3
20	0	0	0	14.5	20	0.3

A convenient formula derived through interpolation for converting the coded variable, x_i , is given as follows (Box and Hunter, 1957):

$$X_i = \frac{X_{iH} - (X_{iH} - X_{iL})(x_{iH} - x_i)}{(x_{iH} - x_{iL})}, \quad i=1, 2, 3 \quad (2.15)$$

where X_{iL} and X_{iH} are the actual low and high levels of X_i , respectively, x_{iL} and x_{iH} are the coded low and high levels of X_i , respectively, and x_i is the coded variable that is being converted.

After the regression coefficients have been obtained, the adequacies of the models are checked using the analysis of variance (ANOVA) technique (Khuri and Cornell, 1987). ANOVA is a statistical method that can be used to quantify the significant differences between factors and levels. It compares the magnitude of the estimated effects of factors with the magnitude of experimental error for the detailed procedure for model validation (Keller, 2001).

2.4.4 Determining the Coordinates of the Stationary Point

After the fitted model is checked for adequacy of fit in the region defined by the coordinates of the design and is found to be adequate, the model is then used to locate the coordinates of the stationary point (Khuri and Cornell, 1987; Zhang *et al.*, 2010). To obtain the coordinates of the stationary point, the fitted second order model in k -variables is written in matrix notation as:

$$\hat{y}(x) = \beta_0 + x'b + x'\beta x \quad (2.16)$$

Where by β_0 , b and β are the estimates of the intercept, linear, and second-order coefficients, respectively. In matrix notation,

$$x = \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_k \end{bmatrix} \quad b = \begin{bmatrix} \beta_1 \\ \beta_2 \\ \vdots \\ \beta_k \end{bmatrix}, \quad \beta = \begin{bmatrix} \beta_{11} & \frac{\beta_{12}}{2} & \dots & \frac{\beta_{1k}}{2} \\ & \beta_{22} & & \frac{\beta_{2k}}{2} \\ & & \ddots & \vdots \\ \text{Symmetric} & & & \frac{\beta_{k-1,k}}{2} \\ & & & \beta_{kk} \end{bmatrix} \quad (2.17)$$

The elements of the $k \times 1$ vector b are the estimated coefficients of the first order terms of Equation 2.16, and the elements of the $k \times k$ symmetric matrix β are the estimated coefficients of the second order terms in Equation 2.16. The partial derivatives of $\hat{y}(x)$ with respect to $x_1, x_2, \dots$ and x_k are,

$$\begin{bmatrix} \frac{\partial \hat{y}(x)}{\partial x_1} \\ \frac{\partial \hat{y}(x)}{\partial x_2} \\ \vdots \\ \frac{\partial \hat{y}(x)}{\partial x_k} \end{bmatrix} = b + 2\beta x \quad (2.18)$$

Setting each of the k derivatives equal to zero and solving for the values of x_i , the coordinates of the stationary point are the values of the elements of the $k \times 1$ vectors x_s given by:

$$x_s = \frac{\beta^{-1}b}{2} \quad (2.19)$$

Where β^{-1} is the inverse of matrix β in Equation 2.19.

The predicted response at the stationary point is given as:

$$y_s = \beta_0 + x_s' b + x_s' \beta x_s = \beta_0 + \frac{1}{2} x_s' b \quad (2.20)$$

2.5 Artificial Neural Networks (ANN)

Biological processes in the food industry, specifically fermentation in brewing, are among the most challenging ones to control and optimise, due to the complexity of the biological raw materials and the use of yeast as processing agents (Linko, 1998). Brewing production is classified into two main stages: wort production and fermentation (Trelea *et al.*, 2004). Although fermentation has been well studied, it still presents challenges to brewers, particularly for industrial scale fermentation. Thus the ability to optimise and predict fermentation would be useful (Gopal *et al.*, 1993, Johnson *et al.*, 1998). Artificial Neural Networks have been gaining a lot of popularity due to their ability to solve complex systems. Their wide spread application is due to the following advantages (Hunt *et al.*, 1992, Widrow and Lehr, 1990, Jain and Sigh, 2003, Zealand *et al.*, 1999):

- Artificial neural networks work effectively even with incomplete data sets;
- They can easily be trained using past data from the experiments or system under study;
- They are applicable to multivariable systems; and
- They have the capability to approximate arbitrary nonlinear functions.

2.5.1 Development of Neural Networks

The motivation for the development of neural networks stemmed from the desire to develop an artificial system that could perform intelligent tasks identical to those executed by the human brain. Artificial Neural networks resemble the human brain in that they acquire knowledge through learning, and the knowledge is stored within inter neuron connection strengths known as weights. They can predict, analyse, associate and emulate the connectivity of biological neurons to solve complex problems in the same manner as the human brain (Huang *et al.*, 2010). Since artificial neural networks were originally adopted from the human neural network system, the nervous system will be used here, to illustrate how the neural network functions, with particular emphasis on the sensory neuron and memory as an example.

The way in which humans recall sensory information is through the aid of neural networks, as follows; each time one is exposed to a particular scene, a sensory signal is passed through the synapses section of the neural network. The synapse is the junction between the neurons and is responsible for directing the signal (towards the appropriate neuron) as well as to store information as shown in Figure 2.5.

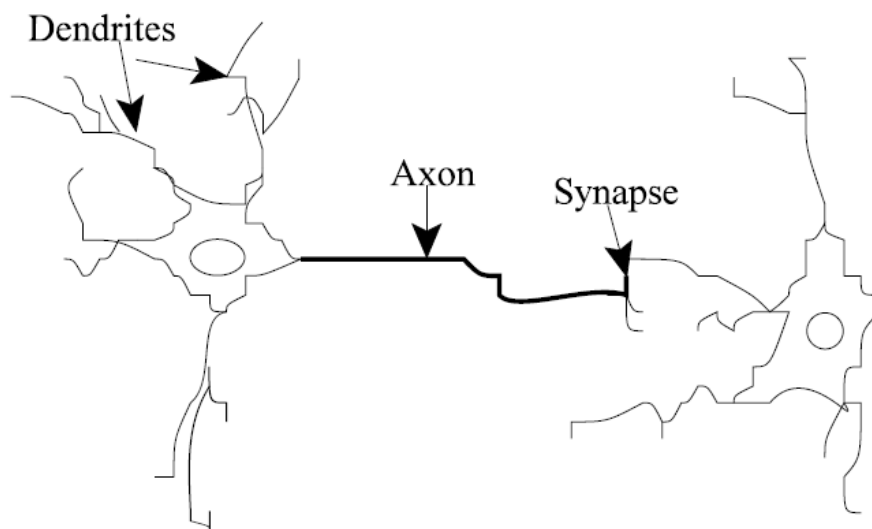


Figure 2.6: Illustration of natural neuron in human.

Natural neurons receive signals through synapses which are located on the dendrites or membrane of the neuron. When the signals received are strong enough such that they surpass certain thresholds, the neuron is activated and emits a signal through the axon. Synapses possess immense powers of becoming more and more capable of transmitting a signal each time the same type of specific sensory signal is passed through it. In medicine, this phenomenon is termed “facilitation” (Guyton and Hall, 2006). Eventually, when the synapse has been exposed to the same type of signal enough times, the signals become so facilitated that the sensations generated in the brain itself (and not necessary felt) will cause impulse transmission through the synapses to the neurons without any real (physical) sensory input. This explains how the human feels what they once felt at the time of the real and an original sensation (sensory signals) even though what they are feeling at the time is only a memory of the sensation (Hanrahan, 2011).

Once these memories are stored in the brain, they are used for future thinking processes; that is to say that the brain will use the stored memories the next time it comes across similar situations upon making responsive decisions. These memories can also be used in the short term memory to facilitate any reflex actions that need to be taken (Beale and Jackson, 1990).

2.5.2 Architecture of the Artificial Neural Network

A similar approach to the description made above of the nervous system is used in Artificial Neural Network programming. Neurons, created on a computer based program, like MATLAB[®], are created that will hopefully perform the same function to those listed above in the human body. The facilitation process described above is more appropriately termed: “training” in Artificial Neural Networks. That is, the exposure of the synapses (and hence neurons) to the same conditions that give rise to certain conditions, time and again until, like the brain, the ANN can transmit the same type of signal the next time it is internally exposed to similar conditions. i.e., given an input, the response generated by the ANN almost perfectly matches what would have been generated in the real life situation (Comrie, 1997).

The most useful neural network in function approximation is the Multi-Layer Perceptron (MLP). Studies by Gardner and Dorling (1997) in atmospheric sciences employed the MLP in a wide variety of tasks, all of which can be

categorised as prediction, function approximation, and pattern classification. The MLP was effective in the prediction and function approximation whereby prediction incorporated the forecasting of future trends in a time series of data given current and previous conditions while the function approximation incorporated modelling the relationships between variables. Predominantly various articles as articulated in Hornik *et al.* (1989) attest that the MLP can be trained to estimate effectively any smooth function and, furthermore, it can model any highly non-linear functions. These features make it strikingly attractive in terms of modelling; hence will focus on its use.

Similar to its biological counterpart in the biological nervous system, as mentioned before, a neural network can learn and therefore be trained to find solutions, recognize patterns, classify data, and forecast future events. In a nutshell, the objective of training is to find an unknown function which relates the input vectors (input data) into output vector (output data) (Bishop, 1995). During training the function is optimised such that the output is much closer to the target value. For a given network architecture, the function is determined by the adjustable weights. The weights resemble memory. Figure 2.6 is a schematic diagram showing an MLP which consists of a single node.

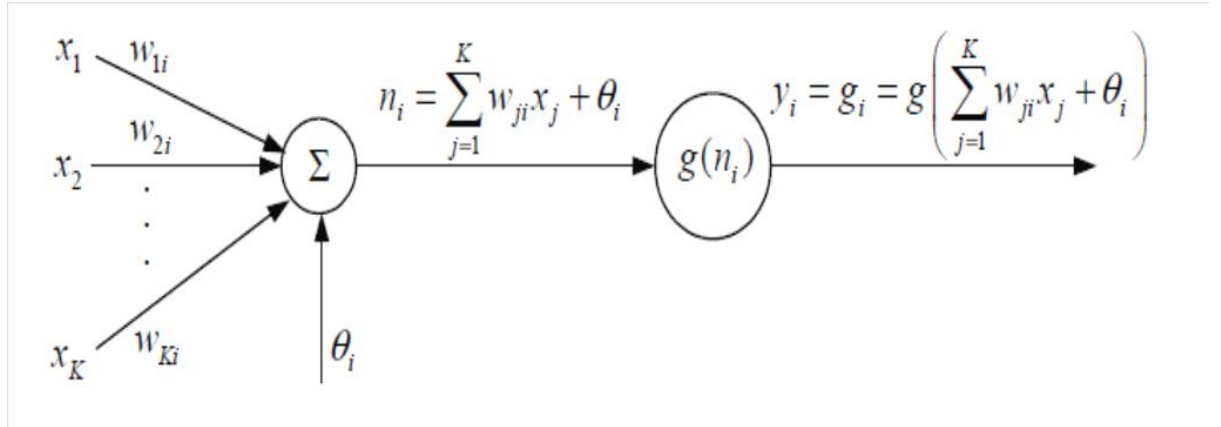


Figure 2.7: Single Node Multi-Layer Perceptron (Patan, 2008)

The steps for computing the output of a single neuron involves calculation of the weighted sum of inputs to the neuron. Thereafter, followed by addition of the bias to the sum and then the sum is fed as input to the activation function of the neuron which then gives the output of the neuron.

Using Figure 2.6, the data inputs, x_k , $k=1,...,k$ are multiplied by weights, w_{ki} and added together with the constant bias term represented by θ_i which give rise to the summation,

$$(2.21)$$

where the summation is the input to the activation function g . From the activation function we obtain the output of the node, which is (Patan, 2008)

$$(2.22)$$

In terms of g which is the activation function, for mathematical convenience, a hyperbolic tangent (tanh) or sigmoid function is mostly used. Hyperbolic tangent is defined as

$$\frac{1 - e^{-2x}}{1 + e^{-2x}} \quad (2.23)$$

A typical multilayer perceptron which consists of several nodes is represented in Figure 2.7.

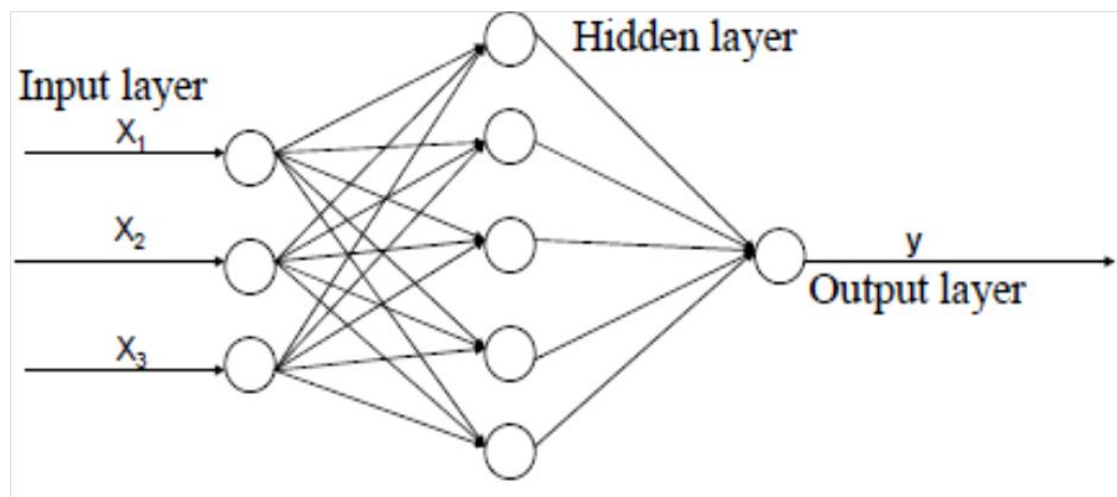


Figure 2.8: Neural Network topology with a single hidden layer (Venkata *et al.*, 2009)

Each neuron in the input layer is connected to a single neuron in the hidden layer and likewise, each neuron in the hidden layer is connected to a single neuron in the output layer to produce the output vector. The information is stored as weights, which are connected between neurons in successive layers and as bias values (neuron activation threshold) (Huang *et al.*, 2007).

2.5.3 Application of Artificial Neural Networks

Artificial neural networks have been used for pattern recognition and data classification in various biological applications such as, modelling the Osmotic dehydration of Mango (Essan, *et al.*, 2009). Zhu *et al.* (2004), studied the optimisation of operating conditions in the rice heat blast process for Chinese rice wine production by ANN and Genetic Algorithms (GA). But little of this type of work has been done for the brewing fermentation process. These include a study made by Bhotmange and Shastri (2011) on the application of ANNs to food and fermentation technology. Pramanik (2004) looked at ANN in predicting batch fermentation variables of ethanol from grapes and *Saccharomyces Cerevisiae* yeast. The main difference is that while Pramanik looked at fermentation in wine production, this study aims to focus on ANN when applied specifically to beer production. Since wine stems from the use of grapes as the source of sugars, beer employs malt/barley for the sugar and starch. Different yeast species are also employed in these two forms of alcohol and therefore, with regard to process conditions, the two ought to be studied separately.

Vassileva and Mileva (2010) produced what may be a more appropriate point of reference to ANN in beer production, Artificial Intelligence (AI-Based) based software tools for beer brewing monitoring and control. In this study, ANNs were used to predict fermentation progress based on the following parameters; properties of the raw materials, *Saccharomyces Uvarum* (*Carlsbergensis*) yeast

conditions and history, as well as fermentation temperature and aeration. In addition to these, the sugar concentration factor is also vital in brewing as it is the source of nourishment for the yeast whose primary function is to convert the sugars into ethanol. It is the sugar concentration that determines the alcohol content of the beer being made and must therefore be studied as well, amongst the factors looked at when studying beer fermentation. This component has been neglected in previous studies and will be included in this study. Further, previous studies placed more emphasis on the use of ANN for process monitoring and control, as opposed to process optimisation. In this study, ANN will be used to model the brewing process after which genetic algorithms will be used to optimise it.

Other important process variables that are considered in this study are fermentation temperature and the yeast concentration. The fermentation temperature is important because it represents the environment in which the yeast cells grow and function (Wu *et al.*, 2009). Some yeast cells will denature while others grow at a specific temperature. For the ale (stout) fermentation, the *Saccharomyces Cerevisiae* yeast species are used while the *Saccharomyces Pastorianus* yeast specie is used in the lager fermentation experiment. Fermentation temperature for the lager yeast should be low (7-8°C) and between 15-22°C for the ale fermentation. Both temperatures are increased by up to 3°C after a couple of days of fermentation. The reason for this is to facilitate the rapid reduction in diacetyl.

Lastly, the yeast cell concentration is an important factor to study. A variation in the yeast may cause variation in fermentation rate or cause off-flavours in the beer being produced (Pollanen and Rousu, 2009). It is for this purpose that the yeast cell concentration was used as a contributing factor in the optimisation of the fermentation process, using Artificial Neural Networks. In addition to this, in their study, Pollanen and Rousu (2009) looked at the yeast factor with regards to fermentation time and speed. They applied the ANN to keep track of the fermentation process at hand but did not incorporate the other important above-mentioned factors, i.e. sugar concentration and fermentation temperature. This study aims to look at the effect of these three factors holistically when optimising the fermentation process.

2.5.4 Importance of Artificial Neural Networks

In the past decades, a number of research methods have been used effectively for the optimisation of engineering processes. However, because today's processes encompass such a wide spectrum, no one optimisation method has proven satisfactory by itself (Bramlette and Bouchard, 1991). Recent studies investigated the use of neural networks and genetic algorithms to model and optimise the production of gluconic acid from glucose (Cheema *et al.*, 2002). However, no incorporation of RSM was made as the experiments were not based on statistical design. Liu *et al.* (1999) found that neural networks outperformed quadratic polynomials in the modelling of a fermentation process. However, the neural

networks were not used in the optimisation step. In this research, study of the use of neural network is reported jointly with genetic algorithms whereby experimental data was also generated using statistical method (RSM) in the optimisation of brewing fermentation processes.

Response surface methodology (RSM) is a method that combines experimental designs and statistical techniques, for empirical model building and optimisation. By conducting experiments and applying regression analysis, RSM seeks to relate a response to some input variables. Challenges have been encountered when the response is complicated; such as consisting of many independent process variables or involving multiple local optimal points. It becomes very difficult to locate the global optimal points because using conventional mathematical methods, the solution converge to a local optimal point (Bramlette and Bouchard, 1991). In reality, multi-responses are often monitored simultaneously. With the increase in both the number of responses and the number of independent process variables, the search space becomes larger and more complex. The optimal values of independent process variables may also differ or contradict one another. This results in a complicated search space with multiple local optimal points and constraints which makes it increasingly difficult to apply mathematical methods to search for the near optimal points. Thus a new method, which is able to search the entire space under the constraints thoroughly and not converge to a local optimal point, is necessary (Baker and Grefenstette, 1989, Zhou *et al.*, 2012, Abbasi, 2012).

Scholars have investigated the possibility of integrating RSM with artificial intelligent techniques such as fuzzy set theory, artificial neural networks, simulated annealing and genetic algorithms (GA). In GA, Parsons and Johnson (1997) applied factorial design and RSM to tune the parameters of a GA program for a DNA assembly problem. These efforts focused on using RSM to determine the parameters of the various applications. Attempts to improve the optimal search techniques being used currently in standard RSM models have not been well studied (Cook and Skadron, 2007, Khoo and Chen, 2001).

Once the ANN has been successfully trained, it can be used as a (non-linear) model for the fermentation process as opposed to the RSM model. Optimisation tools can hence be incorporated to optimise the process, which is the purpose of the study. Genetic Algorithms are chosen because they can perform global as opposed to local searches. This is applicable to the problem at hand because no initial conditions are required in global search which is the case with local optimisation in non-linear problems (Zhu *et al.*, 2004).

2.6 Genetic Algorithms

Genetic Algorithm (GA) is a method for solving both constrained and unconstrained optimisation problems. Like ANN, the method is based on natural selection that mimics biological evolution. Genetic algorithms have been used in conjunction with neural networks to solve problems that are suited for standard

optimisation algorithms, including systems that are highly nonlinear, have slow response, objective function is discontinuous, non-differentiable or stochastic. They have been used in three major ways with conjunction with neural network (Whitley, 1995):

- To set the weights in the fixed architecture of the ANN;
- To learn topology of the neural network; and
- To select training data and interpret the output behaviour of neural networks.

2.6.1 Why Genetic Algorithms?

Genetic algorithms, invented by Holland (1970) in the 1970s, are robust search methods for solving optimisation problems by mimicking the process of natural selection and evolution based on Darwin theory of survival of the fittest. This involves a selection of individuals and the evolution of species as well as reproduction mechanisms and the genetic transmissions of characteristics (Houck *et al.*, 1995). As a result of natural mechanisms, new species are originated ousting those that are not adjusted to their environment. Genetic algorithms have to adjust themselves to certain requirements. They are algorithms that, starting from the initial set of probable solutions to a problem, generate a better and better solution set. New solutions are generated with the use of genetic operators simulating natural reproduction processes.

Each element of the domain of solutions (populations) whose number is constant is called a chromosome, its constituent parts are called genes, while genetic

operators bear the names of selection, crossover and mutation (Golberg, 1989; Baishen *et al.*, 2003). It is the searching of the domain by the GA that generates proposed solutions (the individuals) that are later evaluated by the environment. The role of the genetic algorithm environment is played by the so called fitness function which is defined in all points of the domain which is being searched at a given moment. Optimisation problems like scheduling, game playing, cognitive modelling, optimal problems control, adaptive control, wire routing have been successfully solved using Genetic Algorithms (Michalewicz, 1996). The function also measures the degree to which a given chromosome is adjusted to the requirements of the problem. Therefore, the genetic algorithm can be described as capable of (Tsai, 1993):

1. Generating a population of chromosomes (a set of possible solutions to a problem) randomly;
2. Evaluating each chromosome (each possible solution) in the population according to a fitness criterion;
3. Choosing parent chromosomes with higher performance values in evaluation;
4. Discarding lower performing chromosomes of the population to make room for the new chromosomes;
5. Creating new chromosomes by mating parent chromosomes using mutation and crossover and insert them into the population; and
6. Evaluating the new chromosomes; and if the terminating criterion is achieved, stopping and returning the optimal chromosome (the optimal solution); if not going back to step 3.

2.7 Process Integration

Process integration is family of methods for combining several processes to reduce consumption of resources or harmful emissions to the environment. Process integration for energy saving and waste reduction is a subject of increasing interest among plants operating under continuous and batch mode. In brewing processes, which operates under batch mode, cold and hot streams are scarcely simultaneous. This makes it difficult to integrate processes through direct heat exchange. In such cases, process integration can be achieved through rescheduling of operational timelines (Kemp and Deakin, 1989). This section outlines challenges and opportunities in integration of batch processes. Applications of pinch analysis techniques and time-consideration methods to integrate batch processes are discussed in this section.

2.7.1 Process integration for batch processes

If process integration can be applied inter or intra processes, it has the potential to make the operation more profitable and economical. Ahmad and Hui (1991), in their study they showed the potential of transferring energy between processes, whereby steam in one process was used to heat a stream in another process. Integration between processes is also popular not only because it reduces resources, but it also reduces emissions as well. Figure 2.8 shows the energy consumption of the brewing process at different processing stages. This graph includes the energy required for the whole brewing process. Although the graph is

for a specific production line and size, it does indicate the expectations in the brewing process, extensive use of energy and the opportunities available for process integration.

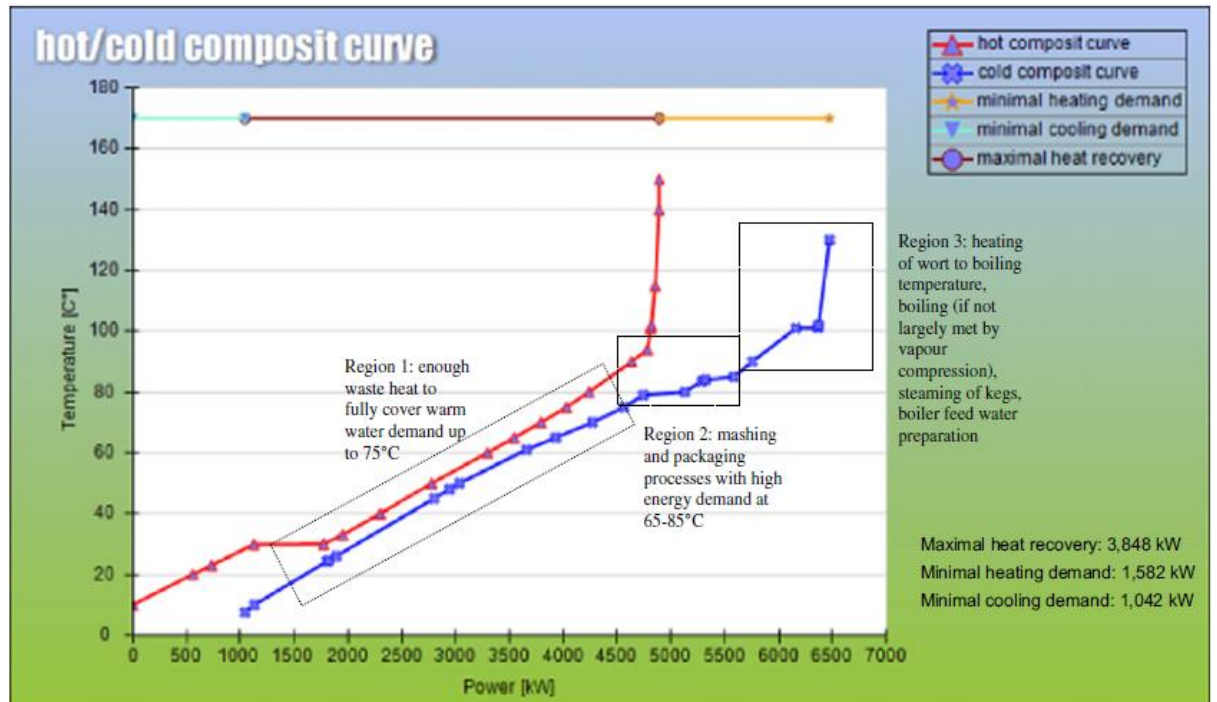


Figure 2.9: Energy consumption curve and possible energy recovery curve in a brewing process (Muster-Slawitsch, 2011)

It was believed in the past that the energy savings for batch processes were less in magnitude compared to continuous processes, hence neglect towards reducing energy utilization in batch plants (Bozan *et al.*, 2001). This belief was proved to be untrue for some batch operations such as dairy products and other biochemical processes (Puigjaner *et al.*, 2007). Batch processes have a wide range of industrial applications and the good instances of these are pharmaceutical, polymer, food and chemical industries (Huang *et al.*, 2007).

Batch process is defined as a series of steps that can be independently operated. Independent operation of the steps involved in batch process allows good control of product quality and yield (Pourali *et al.*, 2006). Batch plants are widely used in chemical process industries relative to other industries (Ashton *et al.*, 1987). They are largely favoured for reactions that take place at a slow rate (Gorsek *et al.*, 1997). In the beer production industry there are different types of processes, continuous processes: semi-batch and batch processes (Priest *et al.*, 2006). Traditional fermentation time ranges from 8-20 days. Most of the brewing processes use batch processing due to the time that is required for the fermentation to occur.

The efficiency of batch processes can be increased by means of process integration where the two most prominent variables considered are the temperature and time (Obeng *et al.*, 1988). This is due to the fact that heat transfer that occurs between the cold and the hot source is not only restricted by temperature but time is involved as well. Temperature constraint clearly specifies that heat can only be transferred between the two streams if and only if a temperature gradient exists between them. The direction of heat transfer is always from the hot stream to a cold stream. On the other hand, the time constraint clearly state that heat can be transferred between the two streams if and only if the two streams co-exist simultaneously within a certain time period (Pourali *et al.*, 2006). Because of the nature of batch production processes, it is rare for matching cold and hot stream to co-exist. Hence rescheduling is of key importance in overcoming temperature and time constraint in batch productions.

Although batch processes are not integrated with ease due to the time constraint; two approaches have been developed that overcome the time constraint problem. The first approach is called shift approach in which there is a transfer of some of the streams from one time interval to another time interval(s). The second approach is storage approach where heat is transferred from a process fluid to an intermediate fluid for storage purposes (Pourali *et al.*, 2006). The stored heat can then be transferred back into the process by matching it with an appropriate stream (Kemp *et al.*, 1989, Sadr-Kazemi *et al.*, 1996).

The combination of the two approaches described leads to other approaches such as Time Average Model (TAM), Time Slice Model (TSM) and those that are based on heat storage (Obeng *et al.*, 1988, Kemp *et al.*, 1989). Furthermore, mathematical techniques such as Mixed Integer Linear Programming (MILP), Mixed Integer Nonlinear Programming (MINLP) have been used to solve the heat integration problem of batch processes (Adonyi *et al.*, 2003). These methods have been applied in test cases involving single-product plants, and heat exchanger networks have been designed for such cases (Zhao *et al.*, 1998). From a practical point of view, there is a clear need for developing new approaches better adapted to the network design of brewing batch plants, and really useful in an industrial environment. With increasing interest in batch process integration, various researchers have suggested various methods to overcome difficulties in batch process integration.

2.7.2 Sequential approach

In this approach, the problem that is being considered is split into two separate parts: scheduling and heat integration (Halim *et al.*, 2009). In the sequential approach, the scheduling part is solved first in order to meet the economic objective such as production time or profit, by using a dedicated optimisation approach. Whereas the heat integration part (this is done to establish minimum utility target) can be solved either through pinch analysis or mathematical optimisation (Lee *et al.*, 1995a, 1995b). The scheduling part considers the availability of a fixed production schedule and applies pinch analysis to reduce the amounts of utilities (Bhatia *et al.*, 1996). In addition, it also assumes a pseudo-continuous behaviour of batch operations. The batch pinch analysis method which was developed during the development of heat integration methods was the heat cascade analysis for heat recovery network design. According to the heat cascade analysis, the process is simply divided into a set of temperature and time intervals within which heat could be cascaded (Kemp *et al.*, 1989a, 1989b, 1989c).

A graphical representation of time pinch analysis was also introduced to solve the heat integration problem; the approach of time pinch analysis involved plotting the energy composite curve in the form of heat transferred as a function of time. The major disadvantage of this method was the involvement of time constraint which limited the amount of heat to be recovered from the hot streams whereas the major advantage of this method lies in its ability to explore different possibilities of cascading through time. In addition, it also allows excess heat to be cascaded to the next temperature interval (Wang *et al.*, 1995).

By comparing the sequential approach with the simultaneous approach, it was found that the key advantage of the sequential approaches lies in their ability to solve heat integration problems without including the simplifying assumptions as opposed to the case when the heat integration problems are solved by using the simultaneous approach (Adonyi *et al.*, 2003). However, the major disadvantage of the sequential approach was the undesirable results of obtaining the suboptimal solutions. The breaking up of sequential methodology into two-step optimisation process does not necessarily imply a poorer solution compared to the simultaneous solution. However, the studies show that it has not received much attention in the literature and there is a great variety of alternate or multiple solutions that occur widely in batch scheduling problems (Kotecha *et al.*, 2003).

Another advantage presented by the sequential method is the fact that the multiple solutions that occur to the scheduling problems allows an individual the flexibility to select a solution that optimizes another (secondary) objective without deterioration in the primary objective function (Wang *et al.*, 1995). For instance, multiple solutions of scheduling allows an individual to choose one solution that requires minimum utilities and the solution chosen may optimize both the primary and secondary objective functions (Kondili *et al.*, 1993).

In addition to the advantages of the sequential method, the sequential method has an ability to tackle problems that are too complex to be solved by using simultaneous optimisation approach (Ierapetritou *et al.*, 1998). For instance most of the problems which were solved using simultaneous method involved the

single objective function (*e.g.*, maximum profit) whereas the problems that were solved by applying the sequential approach involved a number of objective functions, for instance production time, profit, utilities, *etc.* (Halim *et al.*, 2009).

2.7.3 Simultaneous framework

Simultaneous framework is a formulation of an integrated mathematical model comprising of process scheduling and heat integration and solving them both simultaneously for maximizing economic benefits (Mendez *et al.*, 2006). By using the simultaneous approach, the discrete-time scheduling was extended by embedding the heat integration model within the scheduling formulation (Papageorgiou, 1994). The approach is said to require a number of binary variables. It is known that the heat integration of a batch process was neglected due to less recovery, but the problem of task scheduling for maximum energy recovery was later addressed (Lee *et al.*, 1995b). In order to avoid complex formulation, the pairing between the hot and cold streams was previously limited to one-hot-one-cold matching. The one to one heat exchanger rule was later sidelined by applying heat cascade analysis (Zhao, 1998). The result from the simultaneous framework integration was a mixed integer nonlinear programming (MINLP) formulation which can be solved assuming constant heat exchanger time between process streams.

A graphical approach was introduced to solve the scheduling and heat integration problem simultaneously (Adonyi, 2003). The basis of their algorithm was on the branch and bound framework. The branching procedure was used for scheduling of the process units while the bounding procedure was applied for determining the bounds of the utility costs.

A discrete time approach for designing a multipurpose plant with direct heat integration was proposed (Pinto, 2003). In this approach, the heat exchanger network in the form of a superstructure was configured and solved for the optimal structure. This made it a large MILP problem due to the number of binary variables used in its formulation. In addition, there is a proposed continuous model for simultaneous scheduling and direct heat integration (Floudas *et al.*, 2005).

The simultaneous methods are associated with a number of disadvantages. The heat integration problem is over simplified by the method. The availability of a suitable temperature differences as a driving force between heat exchanging streams is generally assumed, than being determined. The energy balance around processing units and streams are neglected and the simple mathematical correlations are generally used to calculate the utility demands needed in the process. The heat exchanger network in the form of superstructures configuration is specified before the utility demands can be calculated (Halim *et al.*, 2009).

These disadvantages have resulted in simultaneous framework being less favoured.

2.7.4 Pinch analysis/ technology

Pinch technology is a systematic methodology for minimizing energy consumption and enhances energy saving opportunities in processes (Linnhoff, 1982). The pinch analysis technique was originally developed for the continuous processes and has to be modified for its application in the integration of a batch process (Corominas *et al.*, 1994). It is important to note that the method has previously been applied with great success for the continuous processes through the processing industry (Kemp *et al.*, 1989c). The batch processes have a time constraint problem, and it is due to this problem that the pinch analysis method be modified such that the time constraint is addressed when dealing with batch process scheduling. The methodology is based on thermodynamic principles and begins with heat and materials balances. The pinch analysis has the ability to rigorously identify and set the process energy targets prior to the heat exchanger network design and can show how this can be achieved practically (Vaselenak *et al.*, 1986, Kemp *et al.*, 1988, 1989b, Linnhoff, 1982). It also ensures that the targets are achieved during network design and can also be set for the utility loads at various levels. The normal types of utilities for a chemical process are cooling water and steam. The methodology provides a consistent way of energy savings from basic heat and material balance to the basic site of utility system (Linnhoff, 1982). Expanding on the modification of the pinch analysis, the targeting method

is said to be easy to apply in a continuous process but complicated to apply in the batch processes because the process streams exist at different times (Kemp *et al.*, 1989a, 1989b, and 1989c). This time constraint problem is addressed by the time average method and the time slice method (Jung *et al.*, 1994). Another approach which was used to address the time issue is heat storage, which was later said that it is not practically possible because the storage system will be losing heat while waiting for the stream to be heated to exist.

2.7.5 Time consideration methods

Figure 2.9 was used for the purpose of illustrating the way in which the heat flow can be computed by using Time Average Method (TAM) and Time Slice Model.

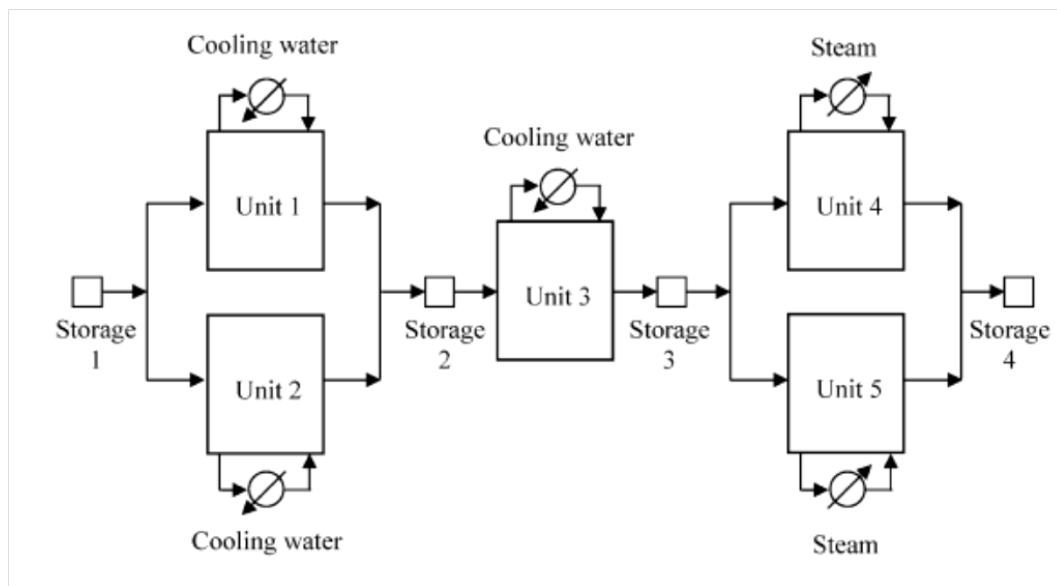


Figure 2.10: Flow sheet for the simple batch process (Halim *et al.*, 2009)

Table 2-4 provides information that is applicable to Figure 2 .9.

Table 2-4 Stream information for simple batch process in figure 2

Unit	Mass(kg)	Status	Start Time(h)	End time(h)	Initial T(°C)	Final T(°C)	Heat Capacity (kJ/kg.°C)
unit 1	100	Hot	1	3	140	60	4.0
unit 4	150	Cold	2	3	40	80	3.0

2.7.5.1 Time Average Model and Time Slice Model

Modifications were made to the heat integration methods that were developed for the continuous processes so that they can be applied to the batch processes. Such modifications led to the model that is now known as time average model (Fernandez *et al.*, 2012). In the Time Average Model (TAM), streams were assumed to be present in all time intervals as if the process was continuous. Time is ignored completely as a constraint and average heat load values are considered over a certain period (Obeng *et al.*, 1988, 2005). The key advantage of this model is the fact that it has been proved to have greater potential for energy integration whereas its major disadvantage is that it was unable to achieve the identified energy targets, when only the direct heat integration was resorted to as a result of ignoring the time constraint (Fernandez *et al.*, 2012). Due to the unjustifiable assumption of ignoring time as a constraint, TAM does not produce accurate results; hence it is referred to as pseudo continuous model. If the pseudo continuous behaviour of the batch process is assumed, then the total heat load of each stream can be determined by averaging it over its availability period. TAM

was also applied to calculate the recovery for waste heat potential (Stolze *et al.*, 1995).

The assumptions made by the TAM were addressed in two ways, namely, time slice model and heat storage (Jung *et al.*, 1994). In the time slice model (TSM), each batch cycle is divided into different time intervals (Ashton *et al.*, 1987). In each of these time intervals, the targets for hot and cold utilities consumption and a local pinch point temperature can be found by using the problem table algorithm. The utility targets of each interval are added together when calculating the minimum utility requirements during a batch period (Halim *et al.*, 2009).

The following calculations show how the heat flows (MJ) can be calculated by using the TAM and the TSM. The calculations were done in conjunction with Table 2.4.

Heat flow calculations

Unit 1: total heat flow = $100 \times 4.0 \times (140 - 60) = 32 \text{ MJ}$

Average heat flow = $32 / (3.0 - 1.0) = 16 \text{ MJ/h}$

Unit 4: total heat flow = $150 \times 3.0 \times (80 - 40) = 18 \text{ MJ}$

Average heat flow = $18 / (3.0 - 2.0) = 18 \text{ MJ/h}$

Figure 2.10 has been used for the purpose of illustrating the difference in the way in which the heat flow can be computed by using the TAM and the TSM

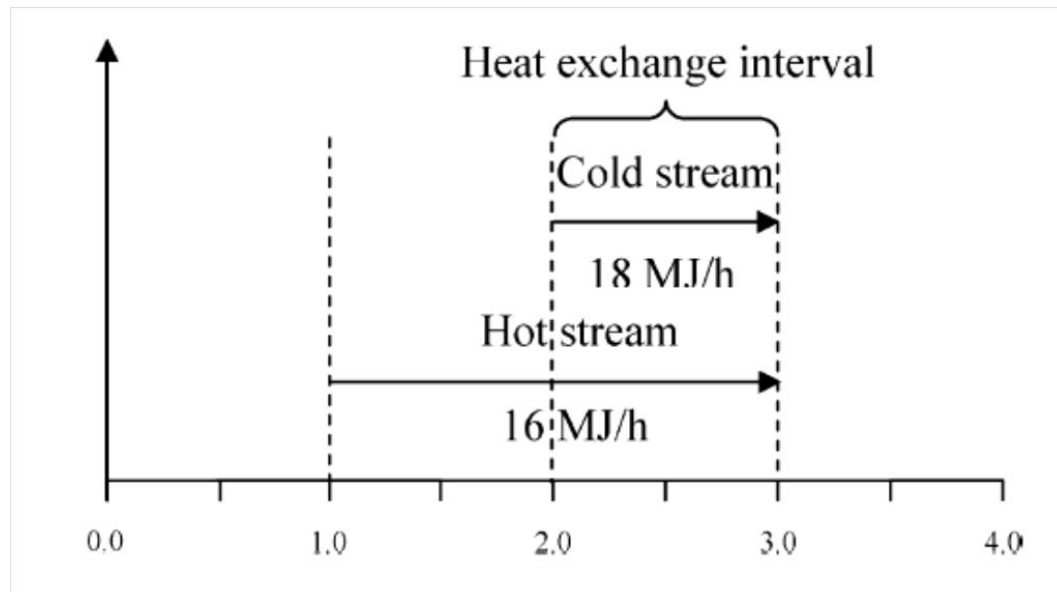


Figure 2.11: TAM and TSM representation (Halim *et al.*, 2009)

From the heat flow calculations it can be noted that, without heat integration, 16 MJ/h of cold utility would be needed for 2 h and 18 MJ/h of hot utility would be needed for 1 h. TSM is applied by splitting the batch period into two time intervals (1-2 h, and 2-3 h) and analyzing the possibility of heat exchange within each (see Figure 2.9). Analysis based on this figure reveals that, in total, a supply of 16 MJ/h of cold utility should be made available during the first interval. In the second interval, however, where the two streams overlap, heat recovery is possible to the extent of 16 MJ. Subsequently, the amount of hot utility to be supplied in this interval is $18 - 16 = 2$ MJ/h. This reduction in required utility indicates the practical applications of TSM and TAM.

2.7.5.3 Mixed Integer Linear Programming

After exploring the various possibilities of heat integration that are offered by scheduling, it was found that a great reduction in utility costs of the batch plants

were made through heat integration combined with the rescheduling of the operational times (Lee *et al.*, 1995a, 1995b). The assumption that was made was that the heat exchange time was negligible compared to the batch processing time and that there were matches between batches of material as they were transferred from one unit to the other. Therefore due to the advantage of great reduction in utility costs, efficient models that lead to the maximum heat recovery were proposed. Hence one of the models that were proposed was the Mixed Integer Linear Programming (MILP) formulation. The model was proposed in an attempt to determine the operating schedule for maximum heat integration between the batch streams for the purpose of reducing the utility consumption (Fernandez *et al.*, 2012). This model focuses mainly on maximizing the heat recovery in batch processes that do not involve the intermediate storage. The main assumption applicable to this model is the cyclical operation of the single product campaigns and one to one matches between the streams. The model has been applied mainly in the large beverage producing plants but it was slightly modified in order to fit the specific industrial conditions (Tokos *et al.*, 2009).

2.7.6.4 Mixed Integer Nonlinear Programming

In addition to the MILP formulation, another approach was proposed which involves a systematic mathematical formulation for the scheduling of batch processes that are operated cyclically. This model assumes no intermediate storage is involved but direct heat integration is possible (Zhao *et al.*, 1998). The model used cascade analysis as a basis and led to what is known as Mixed Integer

Non Linear Programming (MINLP). This model was found to be advantageous as it allows extensive heat recovery from batch processes and it also accommodates the case for multiple stream matching (Fernandez *et al.*, 2012). It was found through this model that an optimum schedule occurs when heat was recovered in the counter current exchange mode but other types of heat exchange can be introduced after rescheduling. The major disadvantage that was encountered in this model was that there are computing difficulties when solving the MINLP model (Fernandez *et al.*, 2012).

Although MILP and MINLP indicate potential energy savings to be realized from batch process integration, they are not practical for a number of reasons. Many streams are available for a particular period of time, which restrict the possibility of heat exchange. Also, for most batch systems, much heating is done *in-situ*; this makes it impossible for them to run at constant temperatures with constant heat capacity flow rates.

In order to incorporate the time constraints in batch process integration, a scheduling approach should be applied. With this approach, streams from the one time interval are transferred to another time interval (s) where process heat exchange between process unit streams can take place.

CHAPTER THREE

Materials and Methods

“It is the weight, not the number of experiments that is to be regarded”

- **Isaac Newton**

3. MATERIALS AND METHODS

3.1 Introduction

In this chapter, preparation of experimental materials used in the study is discussed. Data analysis, experimental and analytical methods are also discussed. A summary is then provided to wrap up the chapter.

3.2 Materials

3.2.1 Wort preparation

Wort preparation was performed in the 100 litre Wits Microbrewery plant. Wort was prepared by first milling 20 kg of malted barley to expose the starch in it and mixing the milled malted barley with water (60 litres) that was heated to 90°C. The pH of the hot water was adjusted by adding phosphoric acid until it came to the required pH of 4. The resulting mixture of milled malted barley and hot acidic water, mash, had a resulting temperature of 70°C. During the preparation of mash, 15 g of calcium chloride (CaCl_2), 15 g of calcium sulphate (CaSO_4) and 15 g of calcium carbonate (CaCO_3) were added to the mixture to act as co-enzymes for the enzymes that would convert starch to sugars during the mashing process. During comparison of batch process to fed-batch, supplements that act as co-enzymes were not added.

The mashing process was then carried out for an hour, after which the sugar solution was removed from the mashing tank. The filtration process in the mash tun was carried out by re-circulating the unclear sugar solution back to the mashing tank where the milled malted barley had by now formed a filter bed to trap most of the small unwanted particles in the sugar solution. The Sugar concentrated solution was then filtered (Scheicher 7 Schell Filter, Germany) to remove the remaining small particles of the milled malted barley that could have gone through the mashing tank sieve. When the sugar solution was clear of malt particles, it was then pumped to the kettle where it was boiled for an hour.

At the start of the boiling process, 2 packets of Brewer's Hops, of 48 g each were added to the boiling solution. After 30 minutes of boiling, a further 2 packets of the same hops were again added. Hops were added during the boiling process for imparting distinctive flavour and aromas to the wort. After the sugar solution had been left to boil a further 30 minutes to complete the one hour of the boiling process, it was then cooled rapidly to give it a heat shock using heat exchanger (BHE Manufacturing Ronneby, Sweden). The hot sugar concentrated solution was cooled from 98°C to 15°C. This cooled sugar concentrated solution is called wort. The experimental set-up for the preparation of wort is presented in Figures 3.1 and 3.2.

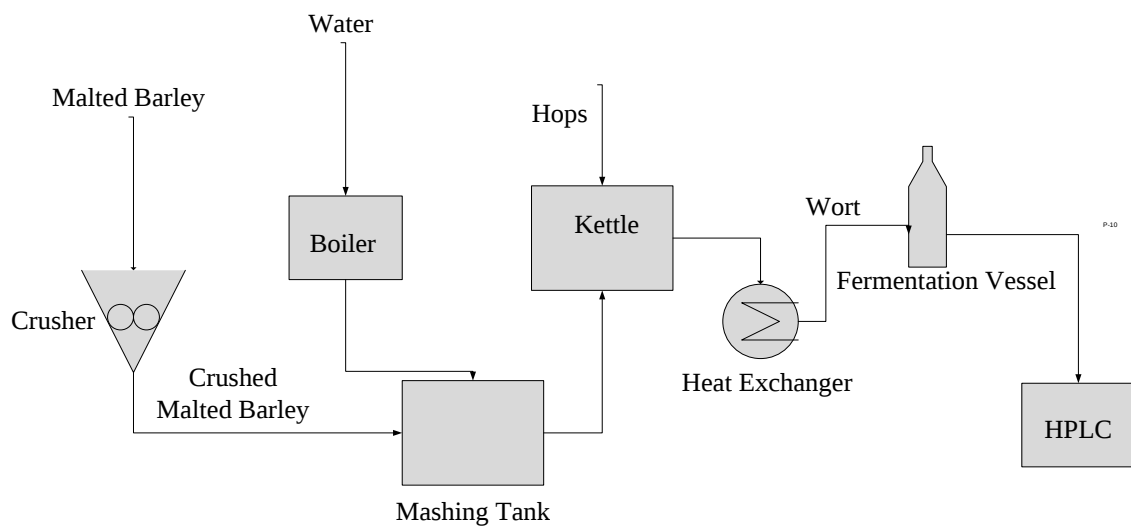


Figure 3.1: Schematic Experimental set-up

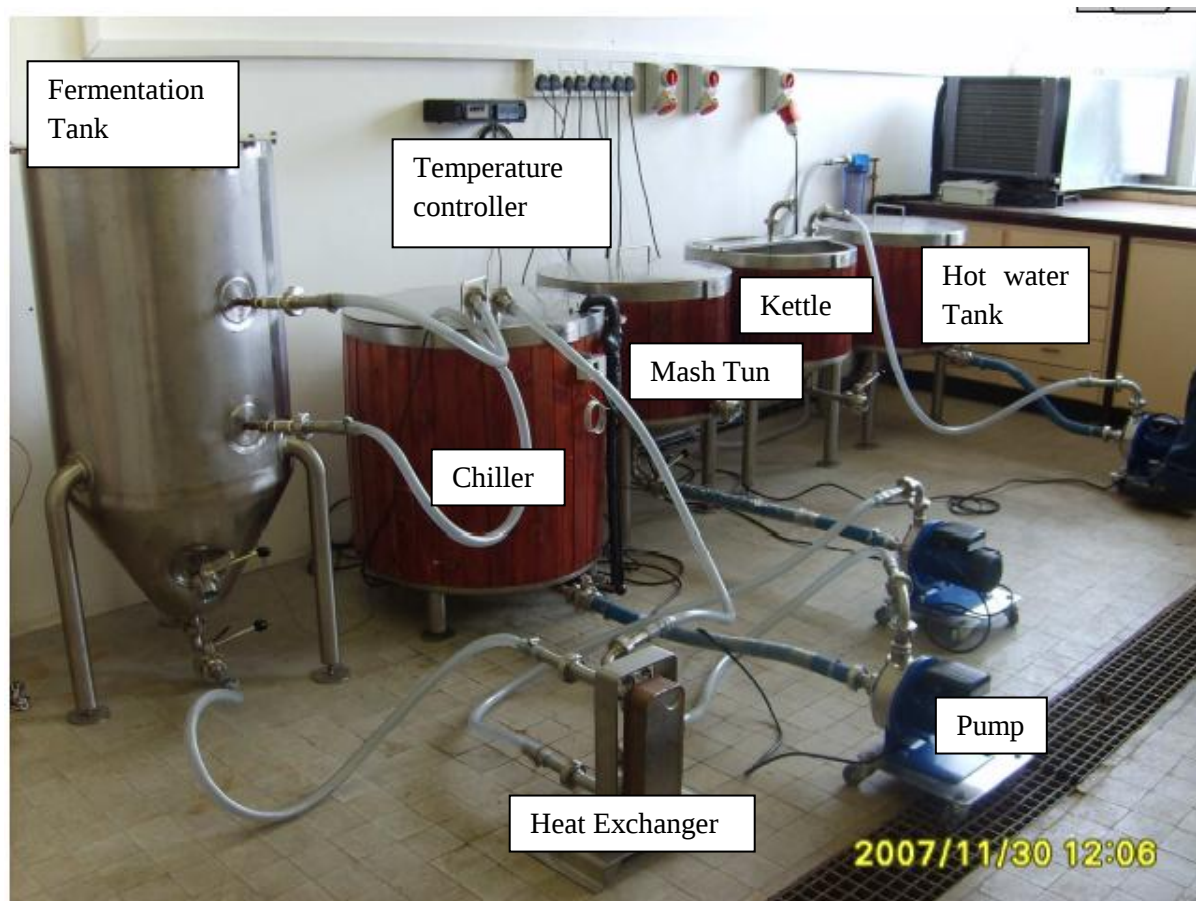


Figure 3.2: Wits mini brewery plant (Madigoe, 2009)

3.2.2 Reagents

All reagents used in this study were of analytical grade obtained from Merck, South Africa (unless otherwise stated), and were used as obtained without any further purification. Drinking water used throughout the experiment was passed through a water filter supplied by Schleicher & Schuell, Dassel, Germany. Barley malt (78% extraction yield, 5.5% moisture) and Saaz Hops (4.5% AAU) were supplied by South African Breweries Ltd-Alrode, South Africa. *Saccharomyces Cerevisiae*, also known as Brewer's yeast, was supplied by National Food Products, Johannesburg, South Africa.

3.2.3 Yeast Preparation

The yeast was propagated in a 250 ml Erlenmeyer shake flask containing 50 ml of 10% (w/w) wort medium. To increase cell concentration, in the second multiplication stage, the yeast was transferred into a 1000 ml Erlenmeyer shake flask with 250 ml of 10% (w/w) wort medium. Pre-cultures were carried out at 30°C and 220 rpm for 24 hours.

3.3 Experimental Methods

3.3.1 Development of Artificial Neural Network model for brewing fermentation system.

Data set for neural network model was collected through multiple set of experiments from Wits Microbrewery and additional data was generated using the

developed RSM model. The data was apportioned as follows: 60% as training set, 20% as validation set and the remaining 20% as testing set. The training set was used to adjust or train the weights in the neural network to produce the desired outcome. The validation data was used to find the best neural-network configuration and training parameters. It was also used to monitor the network error during training. The test set was used only to confirm the actual predictive power of the Neural Network. The criterion used to stop training was the utilisation of regression plot of the validation set and also the plot that compares the output of the NN and its target values. A three layer (input: hidden: output) feed-forward back-propagation NN was used with Levenberg-Marquardt method. Once the ANN model was obtained, it was used to predict the output for the specific set of inputs. However, because the objective was also to find the combination of inputs that can give optimum ethanol yield and rates, different AI technique, Genetic Algorithm, was applied to find the best solution under the specified input. Performance of ANN and RSM were compared using three different sets of experiments.

3.3.2 Reliability of Response Surface Methodology for analysis and optimisation compared to Artificial Neural Networks.

Response Surface Methodology was applied to investigate the interaction of various variables and then used to find the optimum concentration of the main medium components that affect the response, which is the ethanol concentration. A central composite design (CCD) was used in the optimisation of ethanol production. CCD was used for variation of three natural variables that were under

investigation. The experiments were performed in 1Litre fermentation vessels. The solutions prepared were left to ferment for six days.

After testing the fermentation product for ethanol concentration using High Performance Liquid Chromatography, an empirical model was generated using Response Surface Methodology and was used to represent these experimental data graphically so that the optimal yield or production of ethanol can be obtained from the response surface curves.

To check or validate this empirical model, the ANOVA validation method for checking the adequacy of the empirical model was followed.

For the three variables under consideration, a second order polynomial regression model has been proposed by Myers *et al.* (2002) as:

$$y = \beta_0 + \sum_{i=1}^3 \beta_i X_i + \sum_{i=1}^3 \beta_{ii} X_i^2 + \sum_{i=1}^3 \sum_{j=i+1}^3 \beta_{ij} X_i X_j + \varepsilon \quad (3.1)$$

where: y – Predicted response, β_0 – Coefficient of intercept, β_i – Coefficient of linear effect, β_{ii} – Coefficient of quadratic effect, β_{ij} – Coefficient of interaction effect, ε – Term that represents other sources of variability not accounted for by the response function, X_i and X_j – Coded independent variables.

The predictability of the developed model was tested using laboratory experiments.

3.3.3 Comparison of fed-batch fermentation vs. batch fermentation

The main objective of the experiment was to compare ethanol production yield between batch and fed-batch fermentation in order to evaluate the prospects of switching to fed-batch fermentation. The comparison was done without usual addition of substrates that act as co-enzymes in fed-batches. Performances of the fermentations were evaluated at various temperatures. The fed-batch experiments were started with batch fermentations containing 50% of the total batch volume. The additional volume was added at the pre-determined optimal feeding point. The objective function was yield and rate of production of ethanol.

Fermentations of 80 L wort per batch were prepared in a 100 L fermenter at 10°C, 20°C, 30°C and 35°C using the chilling system. Batch fermentation was prepared by adding the entire wort into the fermenter at once. Exponential growth point of batch process was observed to occur between day 1 and day 6. For fed-batch fermentations 40 L of wort was fermented initially for 4 days. Additional wort of 40 L was then introduced into the bioreactor on the same fermentation temperature at day 4. This feeding point was identified as the one to induce an efficient enzyme production. Samples were taken out of the bioreactor once a day and were frozen immediately. At the end of experiment the frozen samples were defrosted, filtered and analyzed. Yeast pitching rate was 1×10^6 viable cells mL⁻¹.

3.3.4 Integration of the brewing process

The concept of rescheduling and time decomposition were applied for batch process integration in order to establish group of streams that have more efficient heat integration. Because of “time-pinch” in batch process, the case study developed was with assumption that multiple batches will be performed at the Micro-brewery. Batch process was initially decomposed into time units and rescheduling was done according to production requirements. Identified time intervals were combined with each other to choose the combination that result in the maximum heat recovery with lower capital investment. When performing process integration, the following assumptions were made:

- No intermediate storage of heat;
- Batch schedule was same for all productions;
- Only hot and cold utilities are considered.

In order to verify the energy improvement of the process, an objective function is evaluated, giving the final energy consumption of hot and cold utilities.

All experiments were carried out in duplicate. Mean values were subjected to analysis of variance (ANOVA) with a *p*value <0.05 using STATGRAPHICS®Plus for Windows 3.0 (Copyright 1994-1997 by Statistical Graphics Corporation).

3.4 Analytical Techniques

3.4.1 Validation of the ANN model

As previously mentioned, validation batches were run in the microbrewery plant to compare the results predicted by ANN model with results obtained from the

experiments. The procedure used is adapted from a similar experimental procedure used earlier. Fermentations for validation of the ANN model were carried out using fermentation system as shown in Figure 3.3.

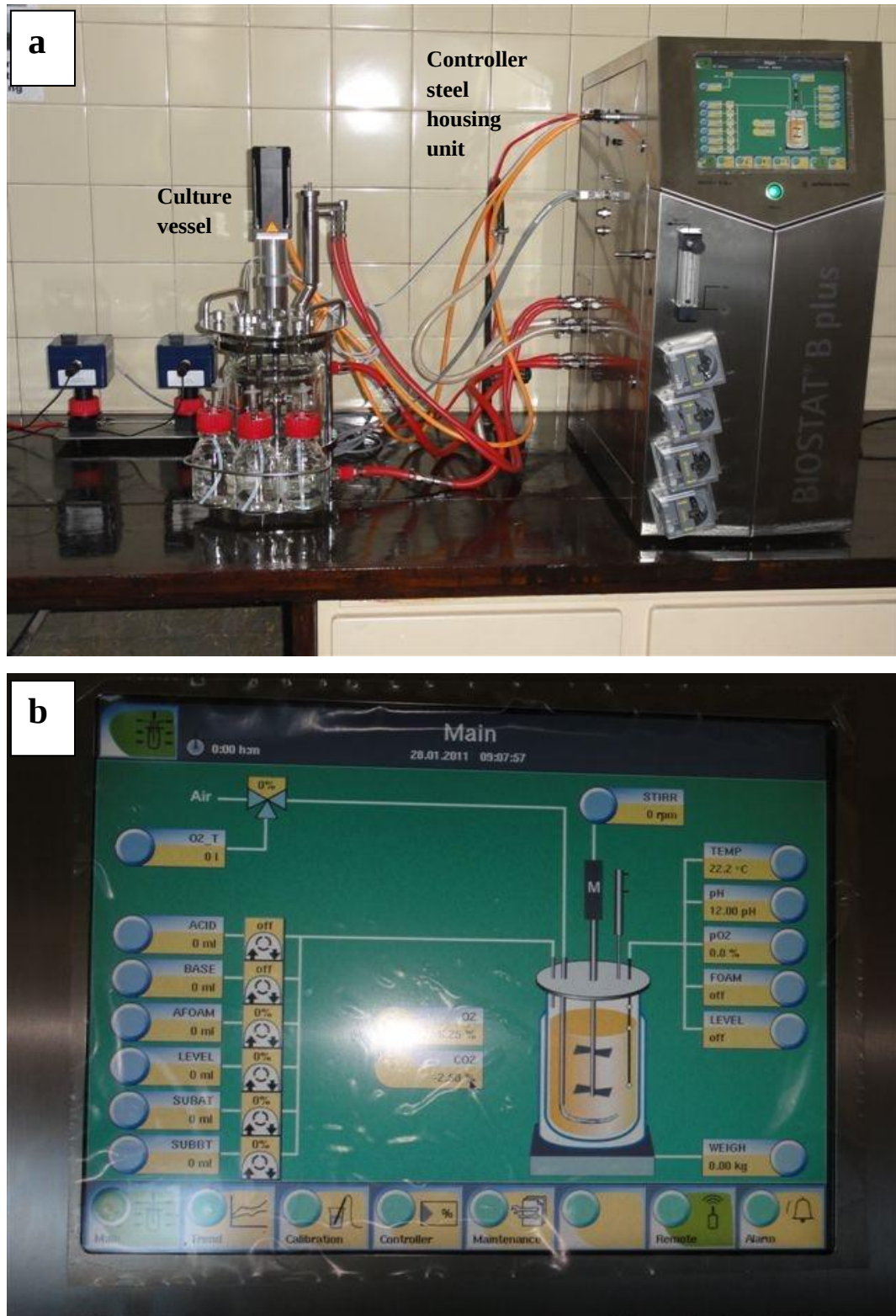


Figure 3.3: (a) BIOSTAT fermentation system used for validation of ANN model, (b) Process flow diagram of the BIOSTAT System

3.4.2 Yeast viability

Yeast viability was determined post-culturing by dilution series and plate counts. The dilution series was carried out using Buffered Peptone Water (BPW, Merck, South Africa) as the diluent. Plate counts were determined using MEA and PDA for *S.Cerevisiae* and the wild yeast species, respectively. Tenfold serial dilution (10^{-1} - 10^{-6}) was carried out in duplicate using 1.5 L broth culture. From each appropriate dilution, 0.1 ml of the diluted suspension was spread plate onto MEA and PDA. Plates were incubated at 30°C for 48 hours. Post-incubation, plates showing between 30 and 300 colonies were selected for enumeration and the pitching rate was determined according to the following equation:-

$$\text{—————} \quad (3.2)$$

Where CFU is the Colony (forming units/ml), N is the number of colonies, D is the dilution factor (dimensionless) and V represents the volume (ml).

3.4.3 Fermentable sugars

Enzymatic and non-enzymatic fermentable sugars were determined by High Performance Liquid Chromatography (Shimadzu USA manufacturing, USA) as shown in Figure 3.5. The HPLC apparatus consisted of Zorbax Carbohydrate Column (4.6 mm ID x 150 mm; 5 μ m), Refractive Index Detector (HP 1100); Guard Column (4.6 mm ID x 12.55 mm) and a 100 μ L Automatic Injector. The Chemstation Software was used for data analysis. Analytical reagent grade chemicals (Sigma Aldrich) were used for HPLC and deionised water used for

preparation of all solutions. Carbohydrate standards (Figure 3.4) used for construction of calibration curve were: - fructose (3.2 mg/mL); glucose (12 mg/mL); sucrose (12 mg/mL); maltose (20 mg/mL); and lactose (20 mg/mL). The sample injection volume was 1 uL.

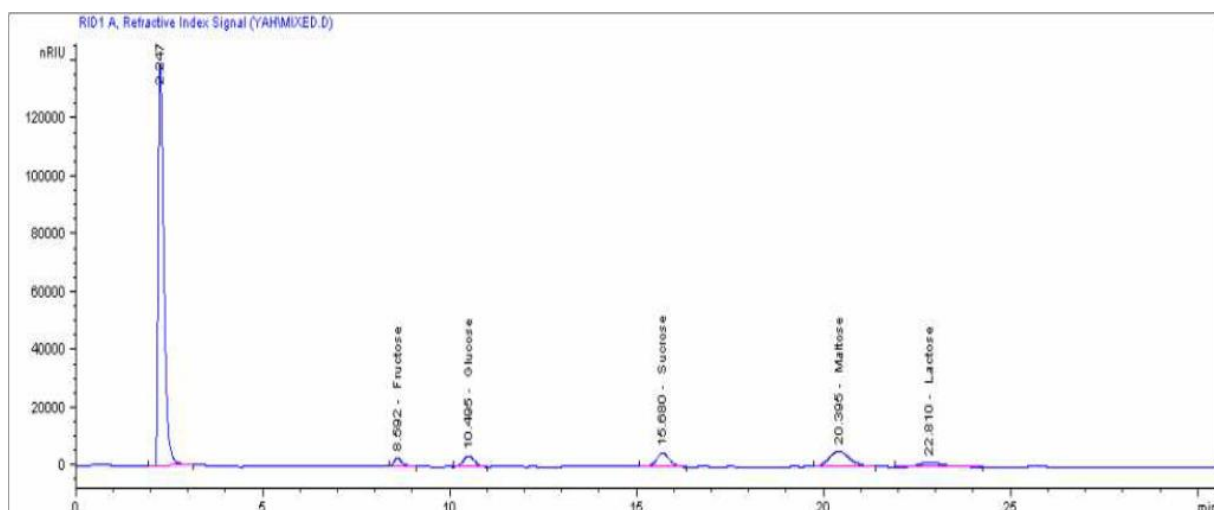


Figure 3.4: HPLC Chromatogram of five Carbohydrate standards.

(Peak 1: Fructose, Peak2: Glucose, Peak 3: Sucrose, Peak 4: Maltose, Peak 5: Lactose)

Post mashing samples were filtered using a 0.22 micron filter (Ultraspin Technology, South Africa) prior to sugar analysis of the wort hydrolyzate by HPLC as shown in Figure 3.5. The system was purged with isopropanol (75:25, v/v) and thereafter standards were run. The column was operated at 30°C, flow rate of 1.4 mL/min, sample injection volume 20 uL and mobile phase of 75/25 Acetonitrile / water (g).

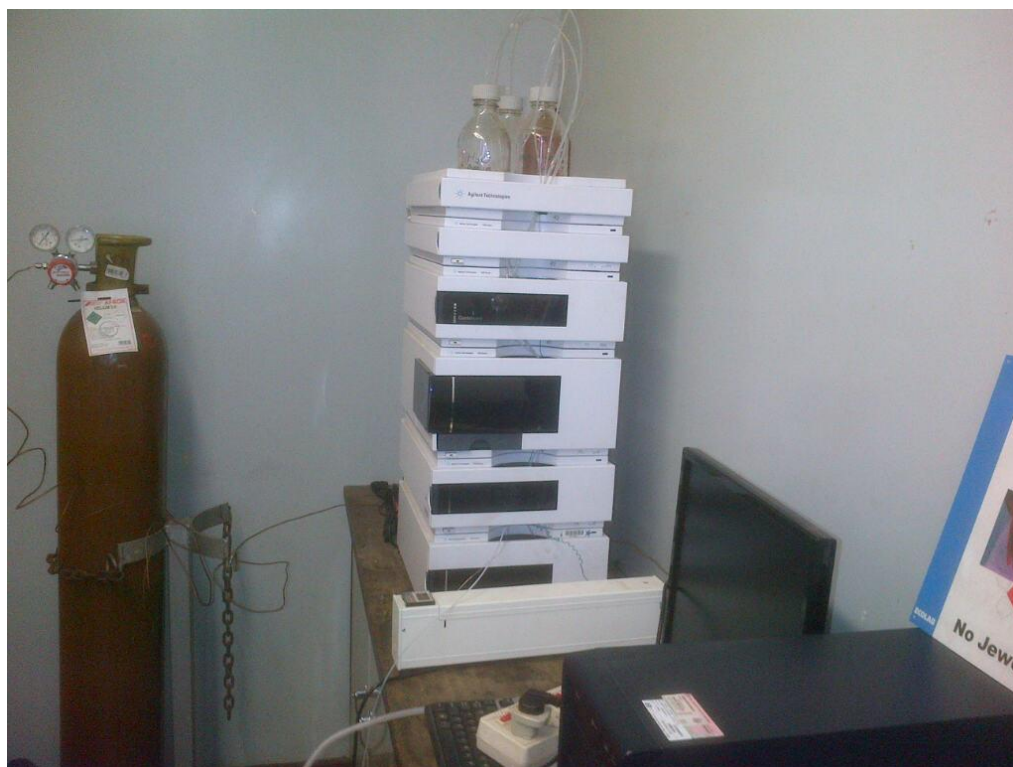


Figure 3.5: High Performance Liquid Chromatography used for sample analysis

3.4.4 Specific gravity

Specific gravity and alcohol content was monitored at 24h intervals using Hydrometer (Wisdom Instruments Factory, China) and Alcoholometer (Wisdom Instruments Factory, China) respectively. Biofuge (Baxter, New York) was used to remove yeast particles from all post-fermentation samples. Five hundred millilitres (500 mL) of the supernatant was used for the reading.

3.4.5 Ethanol content

Fermentation was considered as it was complete when the degree of attenuation reached 75%. The ethanol content was determined according to Gay-Lussac % (v/v). Figure 3.6 represent a standard curve constructed of standard mixtures of absolute ethanol prepared using deionised water used for ethanol analysis.

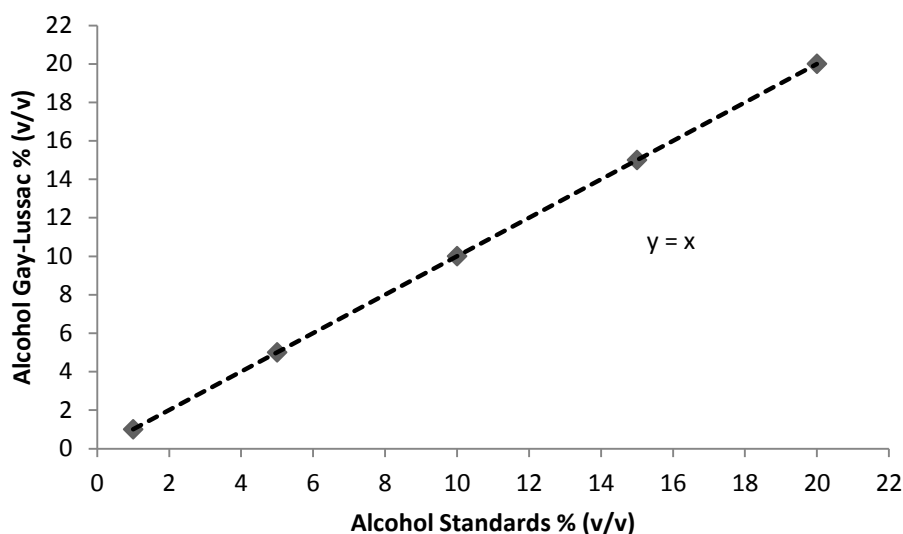


Figure 3. 6: The Alcohol Gay-Lussac % (v/v) of standard mixtures of alcohol.

Readings from the Alkoholometer were validated by using standard sugars. Sugar standards of five carbohydrates (glucose, fructose, maltose, sucrose and lactose) were prepared. For each of five sugars, a fixed volume of alcohol (5%, v/v) was aliquot into 5 g, 10 g, 15 g, and 20 g of sugars and dissolved with deionized water to a final volume of 500 mL. Alkoholometer was then used to measure alcohol percentage for each concentration. Figure 3.6 shows the relationship between measured alcohol and degrees of Gay-Lussac. Degree of Gay-Lussac gives

measurements of alcohol in an alcoholic beverage and is measured by standard Alcohol by Volume as a percentage of alcohol.

3.5 Data Analysis

The data were collected in accordance with the methods described in the previous sections and then processed in response to the investigations posed in the study. Fundamental objectives of the thesis drove the collection of data and the subsequent data analysis. These objectives were achieved. The findings are presented in chapters 4, 5, 6 and 7.

3.6 Summary

This chapter discussed materials and methods used for the study. Analytical techniques and data analysis strategies employed were also discussed. Design of experiments, material preparations and data analysis were guided by the aim of the investigations.

CHAPTER FOUR

Optimisation using RSM

“Simplicity is the ultimate sophistication”

- Leonardo da Vinci

4. OPTIMISATION USING RESPONSE SURFACE METHODOLOGY

4.1 Introduction

This chapter was published in *South African Journal of Chemical Engineers*, 2011, vol. 16, 76-103. Permission to use the article in full has been obtained from the publisher. The author was responsible for conceptualization, experimentation, results analysis and writing of the manuscript. The chapter focuses on optimisation of fermentation using Response Surface Methodology.

Concentration of fermentable sugars from malted barley, yeast and temperature are the most critical variables for fermentation. In the brewing context, fermentation is the production of alcohol from sugars through the action of microbes or by the action of enzymes, which are highly sensitive to temperature. Optimisation of these variables will lead to higher ethanol yield, which will in turn bring positive economic effect on the production of beer. However, optimisation of fermentation process is a very challenging task that relies on an understanding of the complex effects of processing inputs on productivity and quality outputs (Adeyemo and Enitan, 2011). The ‘one-at-a-time’ optimisation technique that is mostly used in biotechnology is not only laborious and time consuming, but it also ignores the interaction among process variables (Haaland, 1989).

Response Surface Methodology (RSM) is preferred to other traditional optimisation methods because it takes into account interaction between variables. It utilises factorial design and regression analysis, which helps in building models to study the interaction and select the optimum combinations of variables for a desirable response (De Connick *et al.*, 2000).

Development of the RSM model that incorporate most important and influential input variables in the brewing fermentation will go a long way in helping industry to select variable combinations for their desired results. The most critical part is the model should be able to predict the response with satisfactory accuracy.

4.2 Materials and methods

In this section the effects of yeast concentration, carbohydrates concentration and temperature on ethanol yield are investigated. The study was conducted by varying concentrations of the three variables.

4.2.1 Experimental Design for RSM and CCD

Fermentation vessels that were used in this study were 1 litre plastic containers. The experiment was carried out by varying the variables; temperature, yeast concentration and media composition in each plastic container. Twenty plastic containers were used for this experiment. 500 ml of wort was poured into each of the 20 containers as the basis of the experiment. Five different values of each of the natural variables were used. Using temperature for example; experiments on this natural variable were conducted in a range of 7°C to 22°C, these are the

minimum and maximum temperature at which ethanol for beer making is produced (Iyuke *et al.*, 2008). Having the maximum and minimum temperatures that this experiment would be conducted, three in between temperatures were calculated according to the Table 4-1.

Table 4-1: Relationship between coded and actual values of the variables (Napier-Munn, 2000)

Code	Actual value of a factor
$-\lambda$	
-1	_____
0	_____
+1	_____
$+\lambda$	

and _____ are the maximum and minimum values of the natural variables respectively, _____ for a

CCRD, k =number of factors studied, $-q$ = fraction of the number of factors (where $q=0$ for full factorial design).

The maximum and minimum coded values, $\pm \lambda$, were calculated to be 1.682 and -1.682 . Thus the coded values that are being used for this experiment which are -1.682 , -1 , 0 , 1 and 1.682 corresponds to the temperature values of 7°C , 10°C , 14.5°C , 19°C and 22°C respectively. Similar calculations were done for the yeast concentration where the minimum and maximum yeast concentrations were 0.1

and 0.5 grams per 500 ml of wort. Thus using the same coded values as that for temperature, the actual values for the yeast concentrations were 0.1, 0.2, 0.3, 0.4 and 0.5 grams per 500 ml of wort.

For media composition, the minimum and maximum sugar concentrations that were used are an original wort sugar concentration for minimum and an addition of 40 g of wort concentrate to the 500 ml of wort for a maximum sugar concentration. Using the same Table 4-1 and the same coded values, the 5 actual different sugar concentrations were calculated to be the wort sugar concentration, addition of 10 g, 20 g, 30 g and 40 g of concentrate to the 500 ml solution of the wort.

A variation of the three principal variables in each plastic container was done by following the central composite rotatable design (Myers and Montgomery, 2002) and is shown in the Table 4-2.

Table 4-2: Experimental layout and runs for the central composite rotatable design

Standard Runs	Coded values of variables			Actual levels of variables		
	Temperature	Sugars	Yeast	T(°C)	Sugars(g/500ml)	Yeast (g/500ml)
Factorial Points						
1	-1	-1	-1	10	10	0.2
2	1	-1	-1	19	10	0.2
3	-1	1	-1	10	30	0.2
4	1	1	-1	19	30	0.2
5	-1	-1	1	10	10	0.4
6	1	-1	1	19	10	0.4
7	-1	1	1	10	30	0.4
8	1	1	1	19	30	0.4
Axial Points						
9	-1.682	0	0	7	20	0.3
10	1.682	0	0	22	20	0.3
11	0	-1.682	0	14.5	wort	0.3
12	0	1.682	0	14.5	40	0.3
13	0	0	-1.682	14.5	20	0.1
14	0	0	1.682	14.5	20	0.5
Centre Points						
15	0	0	0	14.5	20	0.3
16	0	0	0	14.5	20	0.3
17	0	0	0	14.5	20	0.3
18	0	0	0	14.5	20	0.3
19	0	0	0	14.5	20	0.3
20	0	0	0	14.5	20	0.3

The six centre point replications were calculated according to the following equation (Khuri and Cornell, 1987).

$$n_c \approx 0.8385 \left(2^{k/2} + 2 \right)^2 - 2^k - 2k \quad (4.1)$$

where: n_c is Replication points

k is the number of factors studied

The temperature and media composition of the solution being fermented was measured by using a thermometer (Labcon, South Africa) and High Performance Liquid Chromatography (HPLC) (Agilent Technologies, Germany), respectively. Whilst the yeast cells concentration was measured by weighing the respective amount and poured into the 500 ml of the sugar solution. To change the media composition, wort concentrate was added to the sugar solution to increase the concentration of sugars in the media solution or wort. For this experiment, the solutions prepared were left to ferment for 6 days at their various conditions.

After testing the fermenters for ethanol concentration using the HPLC and collecting the experimental data, an empirical model was generated using Response Surface Methodology and was used to represent these experimental data graphically so that the optimal yield or production of ethanol can be obtained from the response surface curves. To validate this empirical model, the ANOVA validation method for checking the adequacy of the empirical model was followed. For the three variables under consideration, a second order polynomial regression model has been proposed by Myers *et al.* (2002) as shown in Equation 3.1.

4.3 Results and discussion

4.3.1 Derivation of fitted model

The data discussed in this section refer to Table 4-3. The detailed statistical and mathematical simulations used in the analysis of the data in this optimisation study were performed using MATLAB[®] R2010a. Detailed sample calculations and results from the HPLC are provided in Appendix B and C respectively.

Table 4-3: Experimental Ethanol Yield

Sample	Actual levels of variables			Observed Yield (mg/ml)	Predicted Yield (mg/ml)
	T (°C)	Sugars(g/500ml)	Yeast (g/500ml)		
Factorial Points					
1	10	10	0.2	10.43	16.09
2	19	10	0.2	25.03	27.85
3	10	30	0.2	23.28	34.35
4	19	30	0.2	53.92	47.92
5	10	10	0.4	14.26	22.13
6	19	10	0.4	63.31	56.51
7	10	30	0.4	22.59	21.65
8	19	30	0.4	59.24	57.85
Axial Points					
9	7	20	0.3	28.89	16.31
10	22	20	0.3	48.39	56.64
11	14.5	wort	0.3	31.92	27.72
12	14.5	40	0.3	44.34	44.20
13	14.5	20	0.1	33.32	26.74
14	14.5	20	0.5	37.93	40.16
Centre Points					
15	14.5	20	0.3	29.22	31.16
16	14.5	20	0.3	33.31	31.16
17	14.5	20	0.3	27.36	31.16
18	14.5	20	0.3	32.59	31.16
19	14.5	20	0.3	36.93	31.16
20	14.5	20	0.3	26.82	31.16

The proposed model for $N=20$ actual observations given in Table 4-3, is expressed in matrix notation as follows (Myers and Montgomery, 2002):

(4.2)

where y is a vector of the observations; X is a matrix of the levels of the independent variables; β is a vector of the estimates of the regression coefficients; ε is a vector of random errors.

The system of equations were solved using the method of least squares. The fitted second order model developed from experimental designs was:

(4.3)

Within the limits: $-1.682 \leq x \leq +1.682$; $i = 1, 2, 3$; where are the coded levels of process variables. Temperature ($^{\circ}\text{C}$), sugars (g/ml), and yeast concentration (g/ml) and $\lambda = 1.682$ is the distance of the star points from the centre of the CCRD that gives the limits of the valid region under experimentation.

4.3.2 Checking the adequacy of the developed model

The statistical adequacy of the second order fitted model was evaluated using the Fisher's variance ratio test known as the *F-test* and the Student *t-tests*. The ANOVA for the fitted model is given in Table 4-4. The ANOVA and acceptance of the model was based on 95% confidence level (Bhatia *et al.*, 2009). Detailed calculations of the RSM model validation are provided in Appendix A.

Table 4-4 ANOVA for the fitted model

Variation	Sum of Squares	Degree of Freedom	Mean Sum of Squares	F - test
Regression	3029.18	9	336.58	4.96
Residual	678.99	10	67.9	
Lack of Fit	602.1	5	-	7.83
Pure Error	78.89	5	-	-
Total	3653.49	19	-	-

Table 4-5 Regression coefficients for refitted model

β	Coefficient	Diagonal Matrix	t-Test
β_0	31.1623	0.166	9.28
β_1	11.9891	0.073	5.39
β_2	4.8982	0.073	2.20
β_3	3.9916	0.073	1.79
β_{11}	1.8768	0.069	0.87
β_{22}	1.6961	0.069	0.78
β_{33}	0.8087	0.069	0.37
β_{12}	0.4562	0.125	0.16
β_{13}	5.657	0.125	1.94
β_{23}	-4.6845	0.125	-1.61

The reliability of the predicted model was tested by comparing the predicted and experimental response results. The observed yield from the experiment and the predicted values are shown in Table 4-3. In this case, there was satisfactory correlation ($R^2=0.8142$) between the predicted values and observed values from the experiment as shown in Figure 4.1.

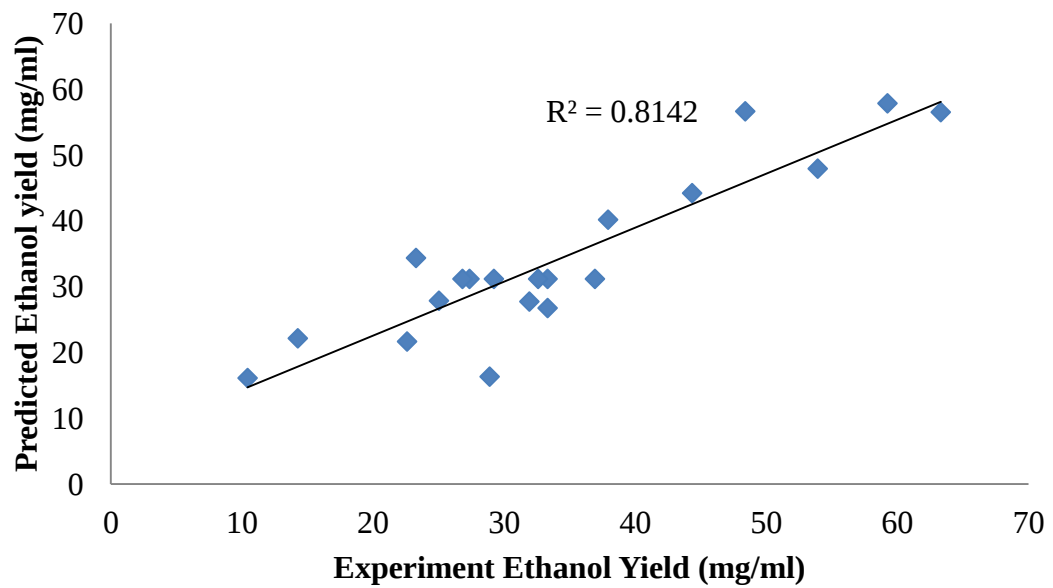


Figure 4.1: Measure of the proportion of the total variation between observed and predicted values.

4.3.3 Response surfaces

Figure 4.2 to 4.3 show the three dimensional response surfaces and their corresponding contour plots of the relationship x_1 and x_3 ; x_1 and x_2 ; x_2 and x_3 when the values x_2 , x_3 and x_1 are held constant at the central level. Thus the resulting equations when x_2 , x_3 and x_1 are kept constant at their central level from the fitted model Equation 4.3 are:

$$y = 31.16 + 11.99x_1 + 3.99x_3 + 1.88x_1^2 + 0.81x_3^2 + 5.06x_1x_3 \quad (4.4)$$

When: $x_2 = 0$, x_1 and $x_3 = -1.682, -1, 0, 1, 1.682$

$$y = 31.16 + 11.99x_1 + 4.90x_2 + 1.88x_1^2 + 1.70x_2^2 + 0.46x_1x_2 \quad (4.5)$$

When: $x_3 = 0$, x_1 and $x_2 = -1.682, -1, 0, 1, 1.682$

$$y = 31.16 + 4.90x_2 + 3.99x_3 + 1.70x_2^2 + 0.81x_3^2 - 4.68x_2x_3 \quad (4.6)$$

When: $x_1 = 0$, x_2 and $x_3 = -1.682, -1, 0, 1, 1.682$

where x_1 – is the coded value for temperature ($^{\circ}\text{C}$), x_2 – is the coded value for sugar concentration (mg/ml), and x_3 – is the coded value for yeast concentration (mg/ml).

4.3.4 Combined effects of temperature and yeast concentration

The combined effect of temperature and yeast concentration on the yield of ethanol was investigated using RSM and the results are shown in Figure 4.2.

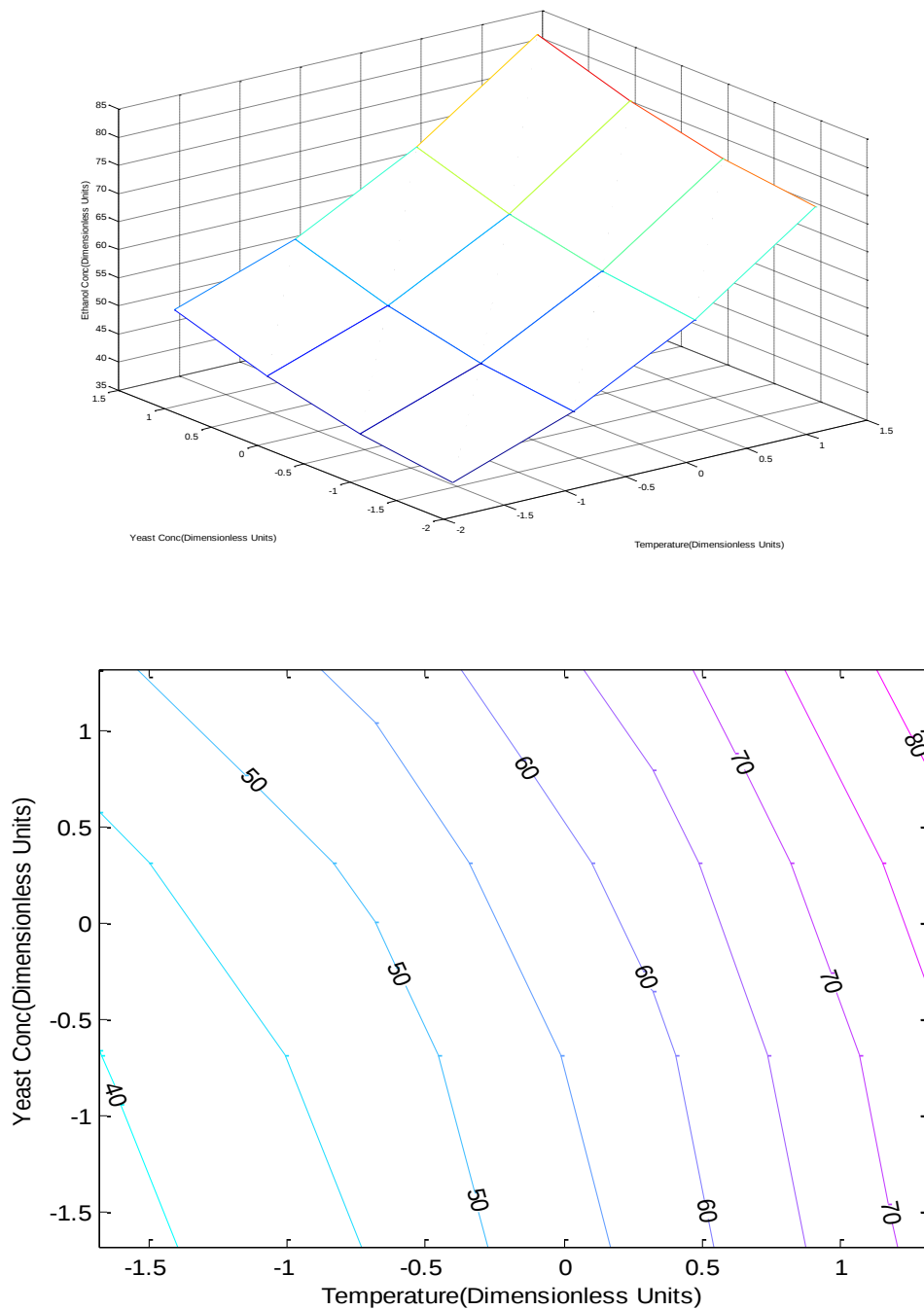


Figure 4.2 Response Surface and contour plots at constant concentration of fermentable sugars

Figure 4.2 shows the effect of temperature and yeast concentration on ethanol yield at constant sugar concentration. It can be seen from the figure that an increase in temperature and yeast concentration at constant sugar concentration increases the yield of ethanol. This can be due to the fact that as yeast

concentration is increased at constant sugar concentrations, there are more yeast cells to convert the sugars to ethanol thus resulting in an increase in ethanol production.

It can also be seen from the figure that an increase in temperature increases the ethanol yield; this can be explained by the fact that an increase in temperature increases the fermentation rate by increasing the rate of yeast metabolism giving higher specific fermentation rates. Thus if a maximum fermentation temperature set-point is specified, then there should be way of removing the excess heat that would be produced during fermentation since this is an exothermic process (Priest and Stewart, 2006). The trend displayed by the predicted model is in agreement with the observed experimental results in Table 4-3, which are consistent with those reported in literature (Ya Ling *et al.*, 1995).

4.3.5 Combined effects of Temperature and sugar concentration

The combined effect of temperature and sugars in fermentation yield was analysed and it was found that both parameters have an effect on the ethanol yield.

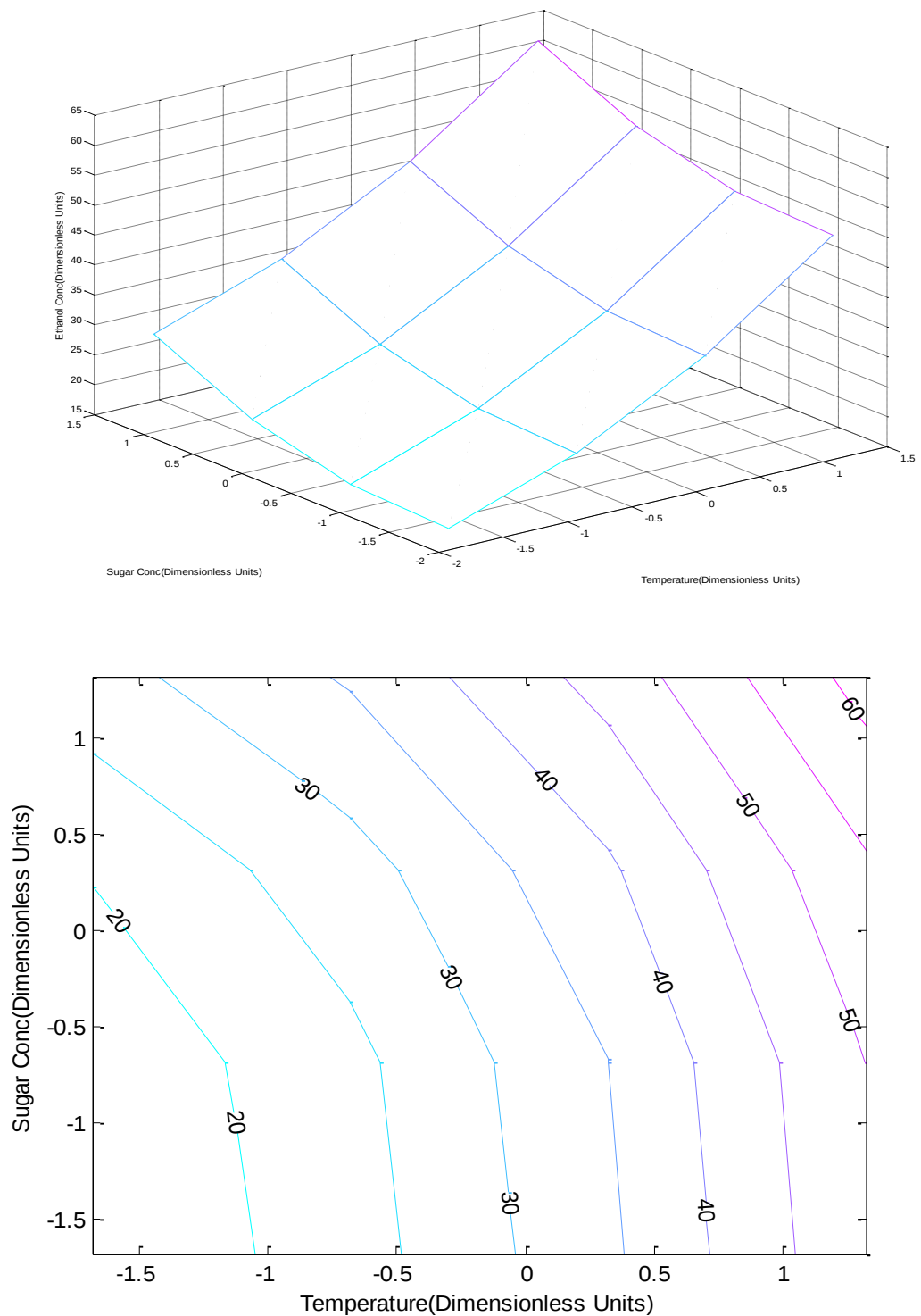


Figure 4.3 Response Surface and contour plots at constant yeast concentration

Figure 4.3 shows the effect of varying temperature and sugar concentration at the centre level of yeast concentration. It can be seen from the figure that an increase

in sugar concentration and temperature increases the yield of ethanol at constant yeast concentration. The reason for an increased ethanol yield due to an increase in sugar concentrations can be due to the fact that the higher the sugar levels, the higher the conversion of sugar to ethanol (Ye *et al.*, 1996). The fact being that, with increased sugars in the medium, the catalyst (yeast) at any point in time will be attached to a substrate to catalyse. The reason for an increase in ethanol yield due to temperature is the same as that discussed above where by increasing temperature results in an increase in fermentation rate by increasing the rate of yeast metabolism giving higher specific fermentation rates (Priest and Stewart, 2006; Phisalaphong *et al.*, 2006).

4.3.6 Combined effects of sugar and yeast concentration

The relationship between sugar and yeast concentration at constant temperature was studied and Figure 4.4 shows the resulting RSM from the model.

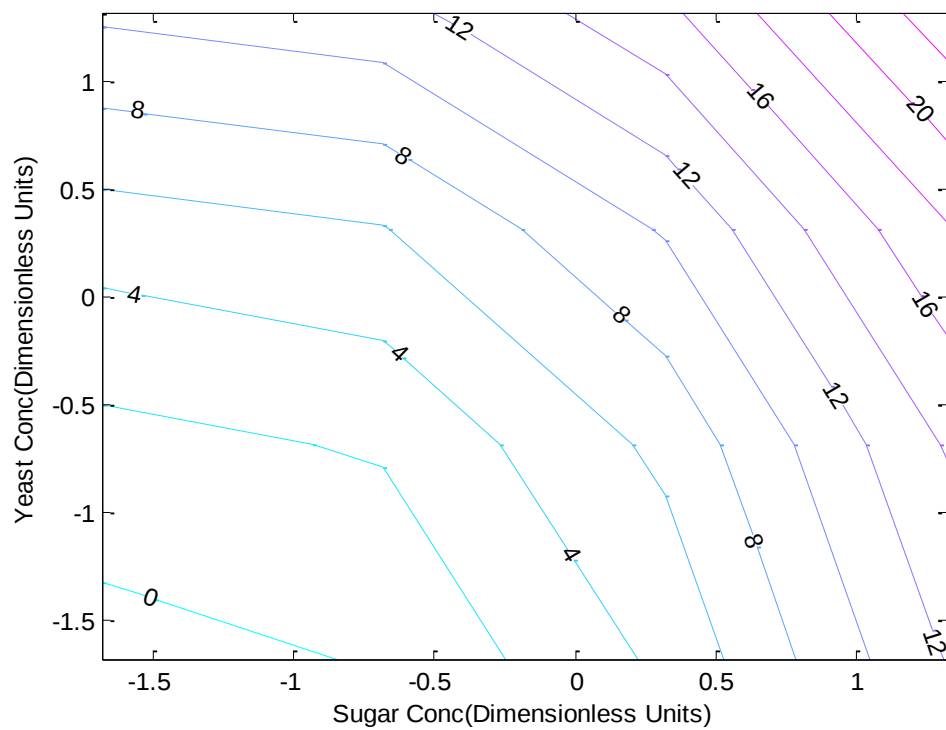
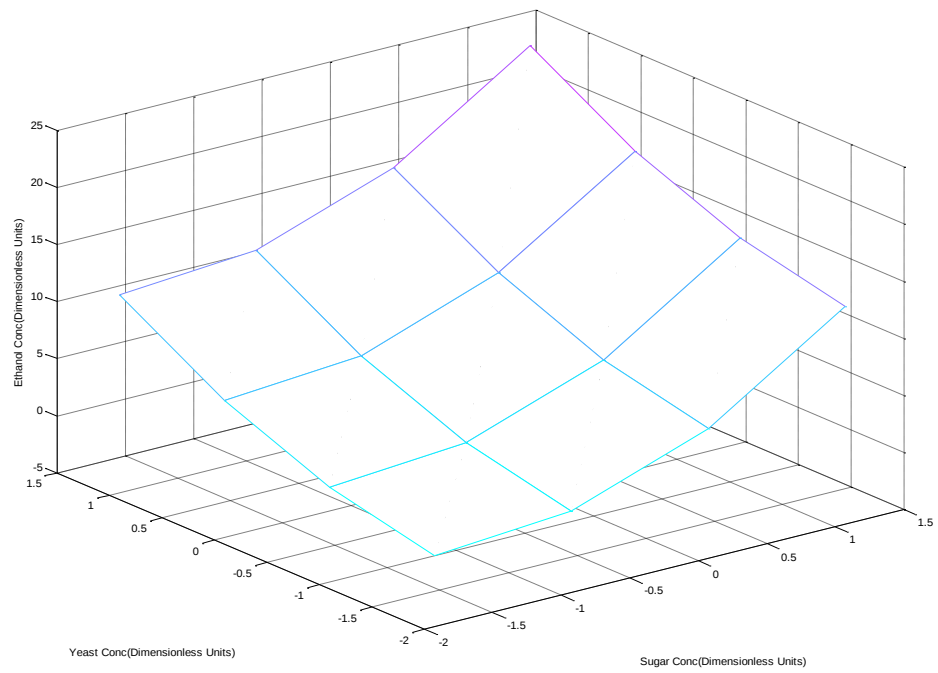


Figure 4.4 Response Surface and contour plots at constant temperature

Figure 4.4 shows that an increase in both, sugar concentrations and yeast concentrations at constant temperature results in an increase in the ethanol yield. This can be explained by the fact that any increase in either yeast concentrations or sugar levels in the fermentation results in more sugars being converted to ethanol, thus increasing the ethanol yield. Increased yield is expected because there are increased sugars to be converted to ethanol, and also there is increased yeast (catalyst) to convert the sugars (Priest and Stewart, 2006; Phisalaphong *et al.*, 2006; Powell *et al.*, 2003).

It can be noted by comparing Figure 4.3 to Figure 4.4 that temperature has the greatest effect on the production of ethanol as it can be seen that when temperature was kept constant at its centre level, and varying sugar concentrations and yeast concentrations this resulted in the least production of ethanol that was produced as seen from the contour plots.

4.3.6 Determination of Optimum Conditions

The developed model was checked for its adequacy and it was found to be adequate, and it was then used to locate the coordinates of the stationery points (Khuri and Cornell, 1987). The conditions of the stationary points were determined using the equation given,

$$x_s = -\frac{\beta^{-1}b}{2} \quad (4.7)$$

, for this study,

$$b = \begin{bmatrix} 11.99 \\ 4.90 \\ 3.99 \end{bmatrix} \text{ and } \beta = \begin{bmatrix} 1.88 & 0 & 0 \\ 0 & 1.70 & 0 \\ 0 & 0 & 0.81 \end{bmatrix}$$

The procedure for obtaining b and β are outlined elsewhere (Myers and Montgomery, 2002; Khuri and Cornell, 1987).

The values of the coded variables were obtained using MATLAB[®] R2010a and resulted in:

$$x = \begin{bmatrix} -3.19 \\ -0.50 \\ -0.50 \end{bmatrix}$$

Since the values are not in the range $-1.682 \leq x \leq 1.682$, they cannot be used. So ridge analysis (Myers and Montgomery, 2002) was used to obtain the stationary points. Thus we have:

$$(\beta - \mu I_k)x = -\frac{b}{2} \quad (4.8)$$

$$\begin{bmatrix} 1.89 - \mu & 0 & 0 \\ 0 & 1.70 - \mu & 0 \\ 0 & 0 & 0.81 - \mu \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \begin{bmatrix} -6.00 \\ -2.45 \\ -2.00 \end{bmatrix}$$

Then

$$\begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \begin{bmatrix} 1.89-\mu & 0 & 0 \\ 0 & 1.70-\mu & 0 \\ 0 & 0 & 0.81-\mu \end{bmatrix}^{-1} \begin{bmatrix} -6.00 \\ -2.45 \\ -2.00 \end{bmatrix}$$

The 3 x 3 matrix is the inverse matrix, it is obtained as follows:

$$A^{-1} = \frac{A^a}{|A|} \quad (4.9)$$

Where A^a is the cofactor matrix and $|A|$ is the determinant of matrix A , hence;

$$\begin{bmatrix} 1.89-\mu & 0 & 0 \\ 0 & 1.70-\mu & 0 \\ 0 & 0 & 0.81-\mu \end{bmatrix}^{-1} = \frac{1}{(2.57-6.07\mu+4.38\mu^2-\mu^3)} \begin{bmatrix} (1.70-\mu)(0.81-\mu) & 0 & 0 \\ 0 & (1.89-\mu)(0.81-\mu) & 0 \\ 0 & 0 & (1.89-\mu)(1.70-\mu) \end{bmatrix}$$

Therefore, the values of the coded variables are,

$$x_1 = \frac{-6.00(1.70-\mu)(0.81-\mu)}{(2.57-6.07\mu+4.38\mu^2-\mu^3)}$$

$$x_2 = \frac{-2.45(1.89-\mu)(0.81-\mu)}{(2.57-6.07\mu+4.38\mu^2-\mu^3)}$$

$$x_3 = \frac{-2.00(1.89-\mu)(1.70-\mu)}{(2.57-6.07\mu+4.38\mu^2-\mu^3)}$$

Since the objective of this study was to determine the conditions that maximise the production of ethanol. Therefore, values of μ greater than 1.70 (the largest of the eigenvalues or characteristic roots) should be used. For this study, a value of 5.6 for μ was substituted to obtain the coded variable. The values obtained are;

$x_1=1.61$, $x_2=0.60$ and $x_3=0.40$, these values are now between the range $-1.682 \leq x \leq 1.682$. It was assumed that the perimeter of the experiment design is at radius of $\sqrt{3}$ units from the design centre point, since the values obtained gave a radius of $\sqrt{3}$.

The optimum yield of ethanol obtained when the optimum coded values are substituted into Equation 4.3 is 63.6 mg/ml. The coded values obtained correspond to a temperature of 21.7°C, sugar concentration of 164.96 mg/ml and yeast concentration of 0.35g per 500 ml of solution; these values were obtained using Equation 2.15. It should be noted that the sugar concentration is a sum of maltose and glucose detected by the HPLC (Agilent Technologies, Germany).

A confirmatory run was carried out under the optimal conditions predicted by the fitted model and it was found that the experiment yield of ethanol under these conditions was 55.8 mg/ml which showed that the model fits the experimental data, with an error margin of 12% because the model predicted an optimal ethanol yield of 63.6 mg/ml. The formulated model was concluded to be satisfactorily valid.

4.15 Summary and conclusions

The main objective in this chapter was to model brewing fermentation using Response Surface Methodology (RSM) with three most important and influential

factors, namely, fermentable sugars concentrations, yeast concentration and temperature. The model could be used for optimisation of brewing process through prediction of the desired response depending on the input variables. The experimental results were analysed statistically guided by the objective of the study.

Although the developed RSM model predicted results acceptable, it has some limitations. One limitation of this method is that the in-between variables (between the maxima and minima) were calculated according to a statistical formula which ensures equal spacing between these values. It thus begs the question, what if one of the conditions that can optimise the system to its absolute optimum lie somewhere between these calculated values that were used? For example, the maxima and minima temperatures were calculated and found to be: 7 °C and 22 °C. After the formula was employed, the in-between temperatures used to run the experiment were calculated as 10 °C and 19 °C. If the experiment ran at say, 15 °C or 16 °C (and keeping the other variables constant), there is a possibility that this would have derived and fitted a better model to the system. It begs the question, what is the basis on which to choose intermediate temperatures to run the experiment? How are these temperatures justified, besides that they are just numerically efficient as they lay strictly in-between the maximum and minimum temperatures and give a good representative temperature for the temperature range at hand?

Studies by Pollanen *et al.* (2001) on neural networks as a tool for fermentation focused on the ways in which to use analytical and computational tools, like ANN, for monitoring the fermentation process, particularly determining the fermentation speed with the hope of early predictions of any problems that might be encountered during the course of fermentation. If these problems can be detected in the early stages of fermentation process, giving enough room and time to re-adjust the system to correct the problems before they actually occur, this would save costs and hence, help optimise the fermentation process. The variables looked at however, are different from this research. They were time from batch start, observed average fermentation speed, specific gravity, time difference between the sample and query time. The raw material properties, yeast condition and history (rather than yeast concentration), tank filling procedure were kept constant. The feed-forward architecture was used for data inputting. Twelve inputs were used in the input layer, and a total of 36 training samples were used. Finally, the back-propagation algorithm (obtained from Neuro Office tool) was used to train the Neural Network.

CHAPTER FIVE

Optimisation using ANN

“Fools ignore complexity. Pragmatists suffer it. Some can avoid it.

Geniuses remove it”

- **Alan Perlis**

5. OPTIMISATION USING ARTIFICIAL NEURAL NETWORKS

5.1 Introduction

Due to the complexity of biological activities associated with fermentation, studies have shown that optimisation methods based on utilising mathematical models and statistically designed experiments are less effective in modelling fermentation processes since they are ineffective in predicting non-linear dynamic systems exhibited by biological processes and thus, there is a great need for new tools that can be employed for the optimisation of fermentation (Vlassides *et al.*, 2000). The importance of this study is reflected on the novel use of Artificial Neural Networks as an optimisation tool for biological systems because of its great potential to circumvent the limitations associated with the traditional optimisation methods used in the fermentation process and thus, it can serve as a basis of discernment in improving the brewing process.

The power of neural networks lies in their ability to represent both linear and non-linear relationships and in addition, they learn these relationships directly from the data being modelled. Furthermore, they rely on inductive inference or generalisation to recognize invariances and capture higher order correlations and generate relationships from inputs and outputs from a set of examples (Becker *et al.*, 2002). The main aim of this study is optimisation of fermentation process variables using artificial neural networks.

5.2 Materials and methods

5.2.1 Modelling the Artificial Neural Network

As indicated earlier in the thesis, the RSM model was revisited whereby the possibility of supplementing RSM with GA's for the determination of near optimal values which in turn were used as inputs in ANN to determine the optimal ethanol concentration were explored. The framework of the prototype is presented in Figure 5.1.

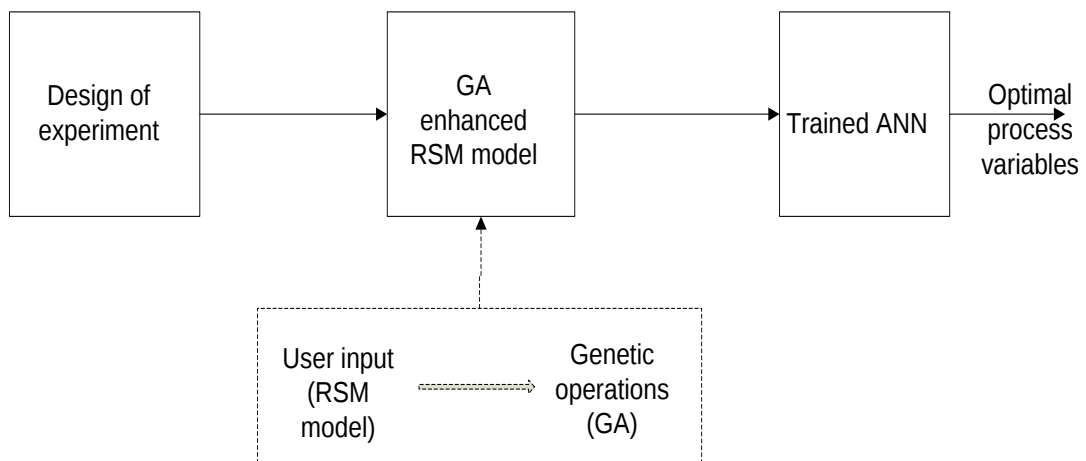


Figure 5.1: Framework of the prototype for intelligent optimisation technique using neural networks and genetic algorithms

5.2.2 Data for artificial neural network (ANN)

Data for the Artificial Neural Network was generated using a model developed using experimental data from the Wits Microbrewery. The governing model

established from the experimental design was Equation 2.15, within the limits: -1.682 for . These limits are the constraints that govern the coded values, that is, (coded value for temperature, °C), (coded value for sugar concentration, mg/ml) and lastly, (coded value for yeast concentration, mg/ml). The principal equations for the coded values were:

$$\text{_____} \quad (5.1)$$

$$\text{_____} \quad (5.2)$$

$$\text{_____} \quad (5.3)$$

For data generation, when choosing the range for the values of it was important to maintain the values of x_i 's between the range: -1.682 . For this study data for the neural network was generated using a value of in the range: 5.7 with incremental steps of 0.1 from the lower limit to yield 24 data points for the Neural Network (NN). The aforementioned range was obtained through trial and error where by the criterion used to check for adequacy was to monitor if the values of the x_i 's were within the limits. The upper limit of can be greater than 8, in this case it was chosen as it would give enough data points for NN training.

5.2.3 Neural network training, validation and testing

The data was divided to produce three independent data sets: a training set, a validation set and a test set. The data was apportioned as follows; 60% as training set, 20% as validation set and the remaining 20% as testing set. The training set was used to adjust or train the weights in the neural network to produce the desired outcome. The validation data was used to find the best neural-network configuration and training parameters. It was also used to monitor the network error during training. The test set was used only to confirm the actual predictive power of the Neural Network. The criterion used to stop training was the utilisation of regression plot of the validation set and also the plot that compares the output of the NN and its target values.

5.2.4 Choosing the Artificial Neural Network Architecture

Concerning the architecture of the Neural Network, the feed-forward neural network was chosen since any transfer function (which maps inputs to outputs) can be used with it unlike its other counterparts (Rumelhart, 1986); and this was advantageous when trying different functions to deduce which one is best. A three layer (input: hidden: output) feed-forward back-propagation NN was used with Levenberg-Marquardt method (method for reducing the error between the projected target and the actual output) which is characterised by guaranteed convergence and numerical robustness (Rajkumar and Bardina, 2005). The default

parameters for Levenberg-Marquardt as indicated in appendix D were used. Another advantage of the Levenberg-Marquardt is the fact that it is 10 to 100 times faster than the typical gradient descent back propagation.

Cybenko (1989) showed that a single hidden layer, given enough neurons can form any mapping needed. In practice, two hidden layers are often used to speed up convergence, while some feed forward networks have been reported in literature to contain as many as five or six hidden layers, the additional layers are not really necessary (McInerney and Dhavan, 1993). The prime thing about NN learning is that it learns best when the mapping is simplest and hence in this study a three layer feed-forward NN was used. The numbers of neurons were varied and the optima were determined by monitoring the prediction of the neural network with the target values and also the error in the validation set. Overall, the best network was selected by identifying the smallest error found with the validation set. Further, the number of iterations was limited to 100 since with greater values; the convergence tends to be extremely slow. During the training phase, the error of tolerance to reach the goal was set at 0.001.

A summary of the procedure employed to select the best configuration of the neural network is summarised as follows:

Step 1- Data was divided into three sessions to form three independent data sets: training set, validation set, and test set.

Step 2-A feed forward neural network was chosen

Step 3-Trained with the training-set data and monitored with the validation set

Step 4-Evaluated the neural network by using the testing data.

Step 5-Repeated steps 2 through 4 with different architectures so to compare responses.

Step 6-The best network was selected by identifying the smallest error found with the validation set.

5.2.5 Determination of optimal values

Neural Networks and Genetic Algorithms (GA) demonstrate powerful problem solving ability. They take advantage of their mathematical nature, non-linear iteration (Koehn, 1994).

In this study, the Genetic Algorithm was integrated as global search method that will determine the optimal value of x_1 and x_2 , that is, the coded value of temperature, sugar concentration and yeast concentration respectively within the desired range (-1.682 to 1.682) and furthermore, use the value of the x_i 's obtained using the GA in the trained neural network to predict the optimal ethanol concentration using the best network architecture. On the other hand, the x_i values were used to determine the actual optimum temperature, sugar and yeast concentration.

The amalgamation of the Genetic Algorithm to the Neural Network for the optimisation is that GA's are widely used as function optimisers (Houck and Joines, 1995). Genetic Algorithms solves both constrained and unconstrained optimisation problems based on a natural selection process that mimics biological evolution. The algorithm repeatedly modifies a population of individual solutions. At each step, GA randomly selects individuals from the current population and uses them as parents to produce the next generation. Over successive generations, the population evolves towards an optimal solution (Nagata and Khim, 2003). The stopping condition to the optimal search is usually dependent on the best fitness that has not changed after a certain number of iterations (Koc *et al.*, 2007). MATLAB provides an optimisation toolbox that includes a GA-solver. The toolbox is executed by typing *optimtool* in the MATLAB's command line and pressing enter. The optimisation window appears and the solver *ga*-Genetic Algorithm is selected and the fitness function (function that needs to be optimised) is inserted by calling the script file. To confirm that indeed over successive generations the population evolved towards an optimal solution, a graph was plotted to confirm convergence.

5.2.5 Checking the adequacy of the developed ANN model

The minimisation of the mean squared error (mse) method was used to assess the performance of the trained neural network model being created. The *epoch* here is the presentation of the input vectors to the network as well as the calculation of

the new weights and biases. The training parameters are presented (*epoch*) one at a time, as opposed to batch training and the training stopped when epoch was on the 21st iteration.

The performance is the behaviour of the network and the performance function (used to analyse the performance of network outputs) was employed. It is also displayed in graphical form in Figure 5.2 which shows a plot of the mean squared error versus the *epoch* (or iteration number).

Three stopping criteria were defined as:

- If the maximum number of *epochs* are reached;
- If the performance goal is met; and
- If any other stopping condition of the training function occurs.

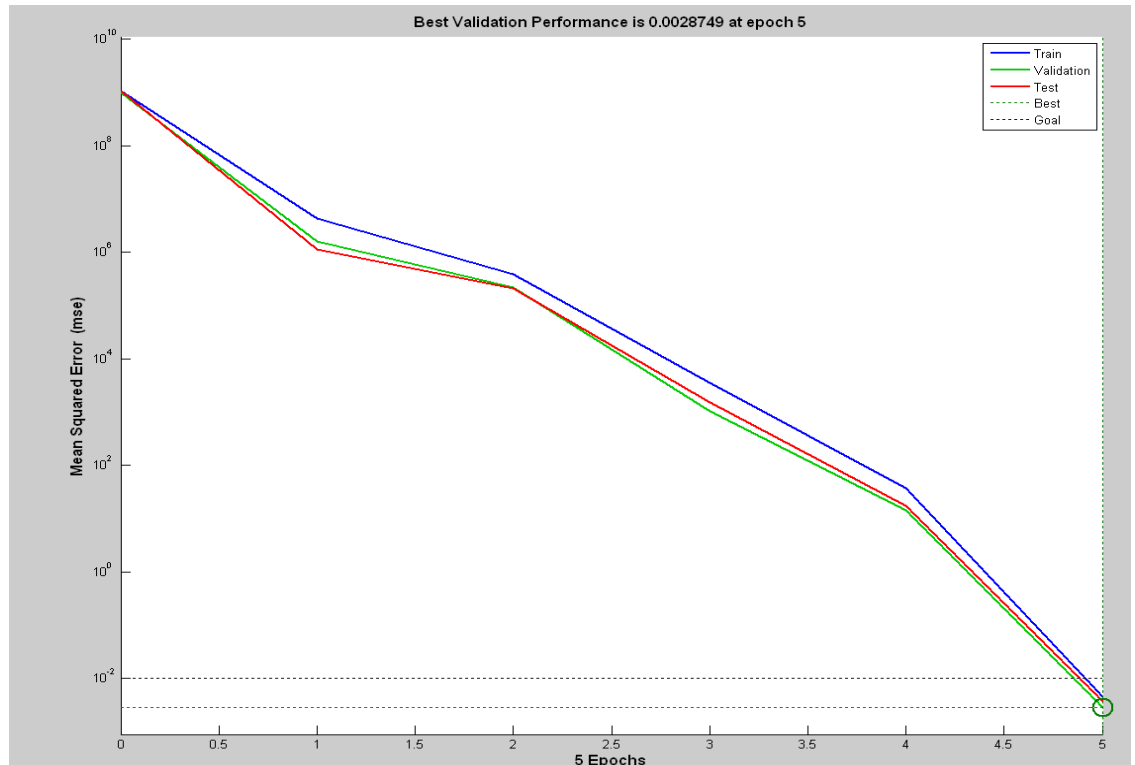


Figure 5.2: Train, Validation and Test data mean squared error against epoch

The gradient descent method was used to minimise the mean squared error (mse). Here, the changes made to the weights, biases and the associated adjustments were made inversely proportional to the derivatives of the neural network error. Here, it can be seen that the gradient is hence allowed to decrease within the range from . Note the minimum gradient () displayed on the training window is an indication of the default settings that come with Levenberg-Marquardt training function displayed in Table D-1 in Appendix D.

μ , is a parameter for the initial value of the scalar, μ . It should not be confused with μ , in the function file called train, in Appendix D that was used to generate

the three coded variables (temperature, sugar and yeast concentration). The MATLAB functions are shown in Appendix D.

Validation checks are made to confirm that the trained network closely approximates the target inputs. Validation checks are used as part of the stopping criteria for training the neural network. The training stops when the validation checks have been made or if the performance goal has been met, which was the case in this study.

5.3 Results and Discussion

5.3.1 Artificial Neural Network (ANN) model

The Neural network training was accomplished by using the training, validation and testing data sets. The number of the hidden neurons is an important design issue and in this study it was taken into consideration in terms of selecting the best network architecture. On the one hand, having more hidden neurons allows the network to approximate functions of greater complexity but, as a result of networks high degree of freedom, it may overfit the training data while the unseen data will be poorly fit to the desired function. On the other hand, although a small network will not have enough power to overfit the training data, it may be too small to adequately represent the target values. In order to choose a reasonable amount of hidden neurons, three different networks with 2, 10 and 30 hidden neurons were examined.

The NN was trained several times to predict the target values. Figures 5.3, 5.4 and 5.5 are the plots showing the target values (obtained from experimental model) and output values (predicted by Neural Network) using 2, 10 and 30 hidden neurons respectively. The same number of training runs was done to each NN and results were compared to select a better architecture.

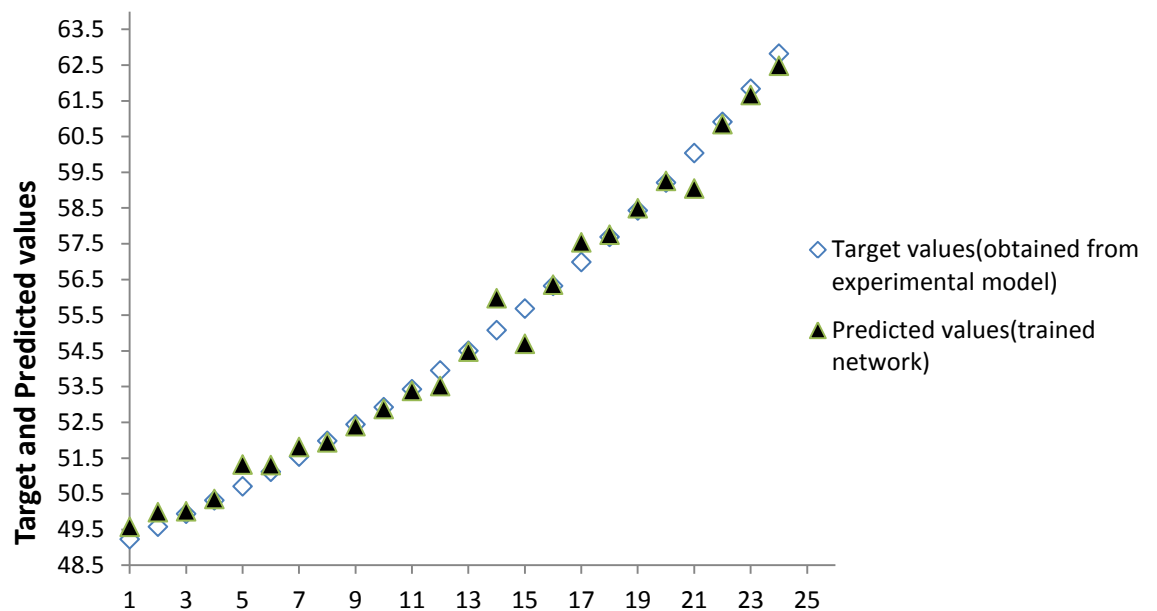


Figure 5.3: Plot of Target and Predicted values when using 2 neurons in hidden layer

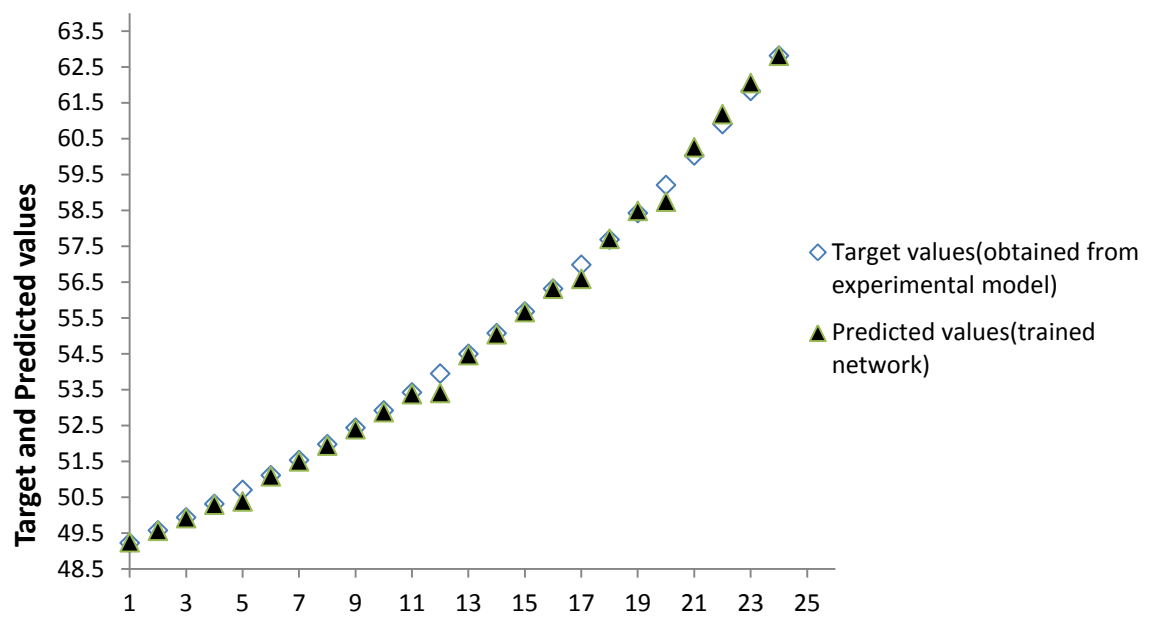


Figure 5.4: Plot of Target and Predicted values when using 10 neurons in hidden layer

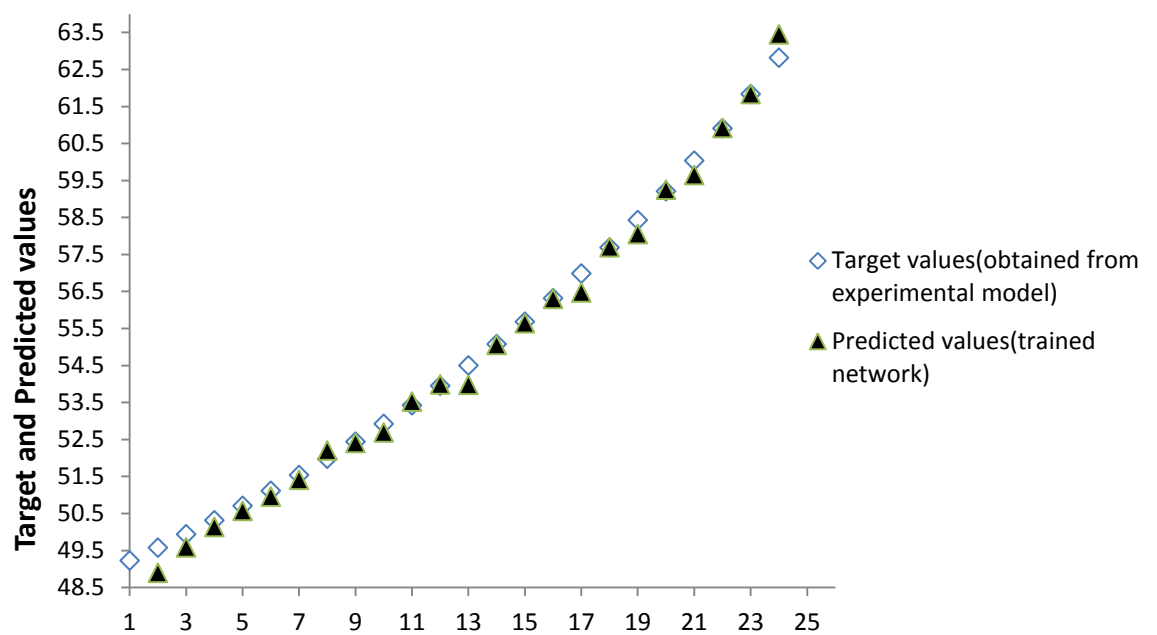


Figure 5.5: Plot of target and predicted values using 30 neurons in hidden layer

From comparison of Figures (5.4 & 5.5), it can be deduced that the network with 10 hidden neurons outperforms the other two networks after they are trained with the same data parameters since the predicted values are much closer to the target values, hence, the Neural Network architecture (3:10:1, that is, 3 inputs, 10 neurons in hidden layer and 1 output which is ethanol concentration) is quite good and it was used to predict the ethanol concentration once the optimal x_i values (coded values of temperature, sugar and yeast concentrations) have been determined with the use of Genetic Algorithm.

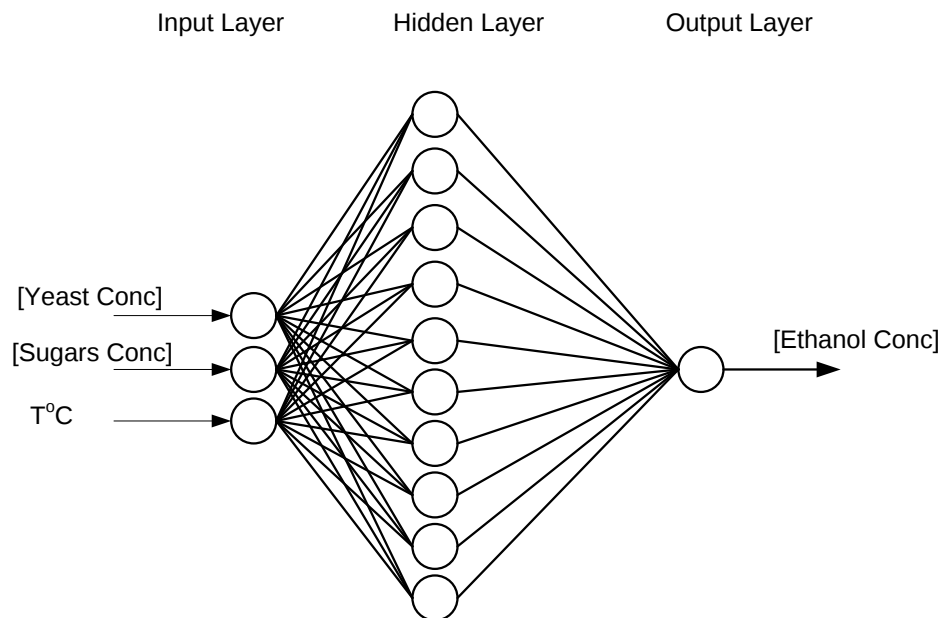


Figure 5.6: Architecture of the developed Neural Network Model

The subsequent validation and testing plots considered in Figure 5.7 and 5.8 were for the network with 10 hidden neurons.

Figure 5.7 shows the Network validation data plotted with comparison to the line $Y=T$ which is the zero error line, that means if data lies on the zero line or the best fit line coincides with the zero line, no error experienced, target values would be identical to predicted values by NN. The circles represent some of the validation data and the solid line is the line of best fit. This graph was generated by the MATLAB code which automates the Neural Network Toolbox that shows the error and regression plots.

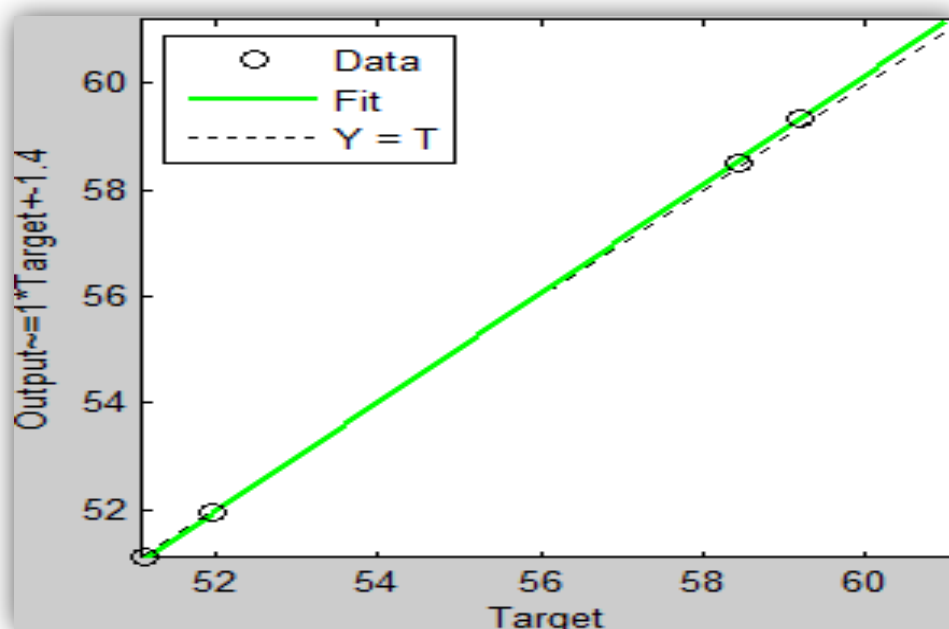


Figure 5.7: Validation plot for ANN model

From Figure 5.7, it can be seen the validation data and the line of fit almost coincide or lie close to the zero error line (dotted line) which attests that the NN has a great potential to predicting data very well. Regression R-value directly shows the network performance ($R^2=1$ means perfect correlation). Several configurations of the neural network were considered and data was trained and the

R-value of the validation curve was monitored to select the best configuration and also to stop training. Further indication of the improved performance of the developed ANN is indicated in Figure D.5 in Appendix D.

From the testing data, Figure 5.8 shows a plot of the test data against the zero error line ($Y=T$). If the training were perfect, again the data or the best fit line would coincide with the zero error line.

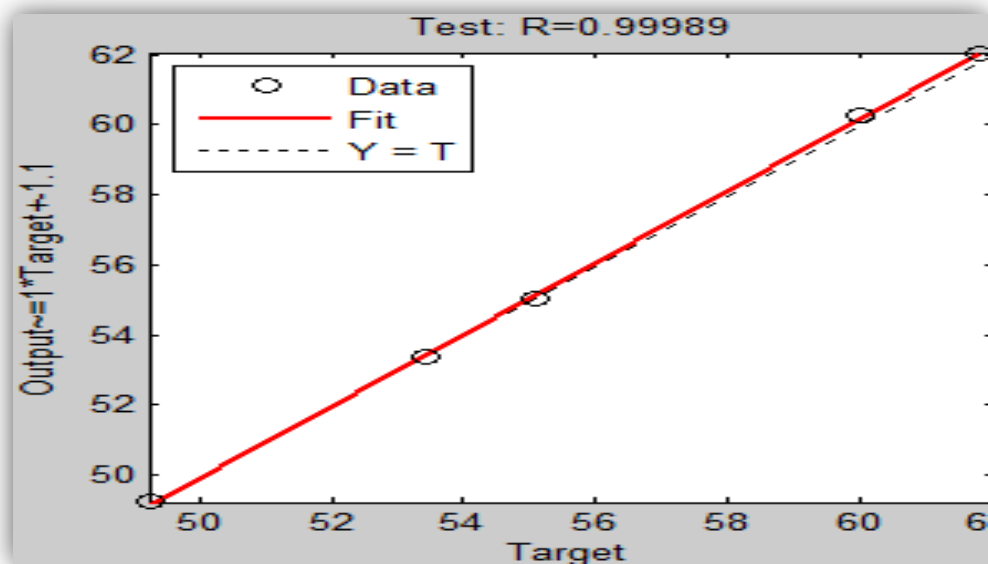


Figure 5.8: Test data for ANN architecture

Likewise as discussed before, the value of R^2 is 0.9989 which infers a good NN architecture and hence from the above figures it was deduced that the network can be used to estimate the relationships between the process performance variables and output with reasonable high degree of accuracy and hence the Genetic

Algorithm was used to find the optimal value of x_i 's and then use them in the trained network to predict the optimal ethanol concentration.

Genetic Algorithm was used for the calculation of the optimal values. They use the three main principles of natural evolution: reproduction, natural selection and diversity of the species, maintained by the differences of each generation with the previous. As shown in Figure 5.9, the optimal values of x_i 's were obtained after 51 generations. The optimal x_i values were 1.0025, 0.9998 and 1.0012 for the coded input values of temperature, sugar and yeast concentrations respectively which falls within the expected limits (-1.682). The objective function value was -45.83, is related to the sum of squared errors, this number can be large, small or even negative. This value depends among other things, on the amount of data itself. The numerical value of the objective function has no meaning by itself (Kalogirou, 2004).

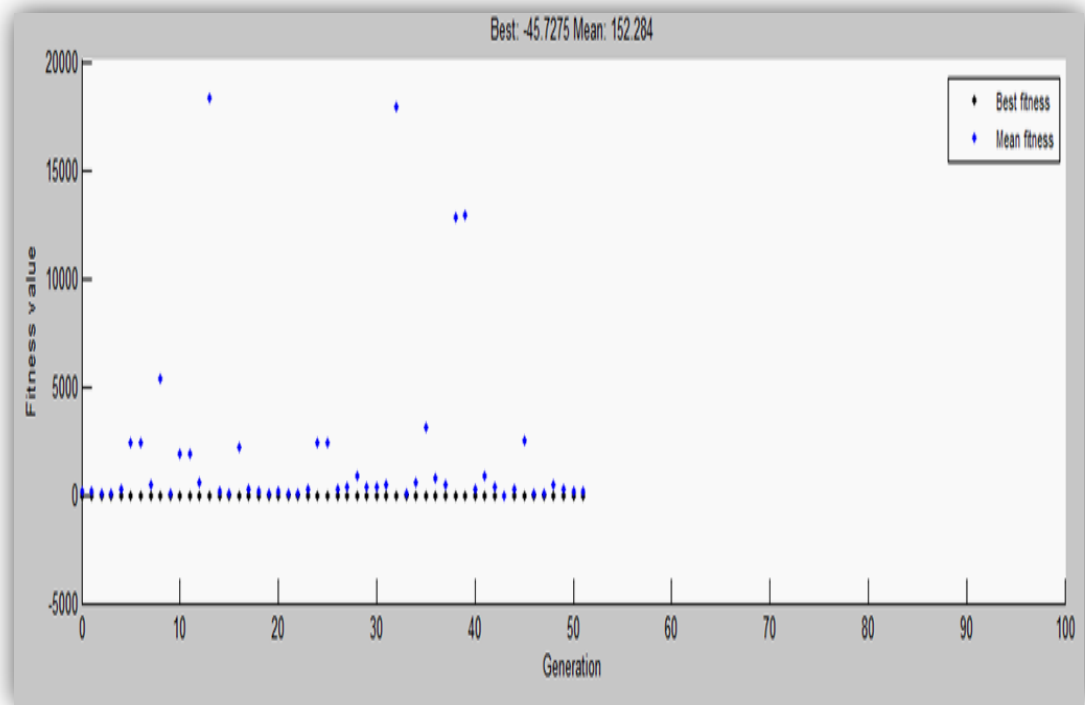


Figure 5.9: Genetic Algorithm population plot

From Figure 5.9, it can be inferred that optimisation took place since it can be observed that the population converged, that is, as one moves from left to right from the plot, the GA randomly selected individuals from the current population and used them as parents to produce the next generation. Over successive generations, the population evolved towards an optimal solution and this was achieved in the 51'st generation.

The optimal values of temperature, sugar concentration and yeast concentration were found to be 19.0°C, 163.9 mg/ml and 0.42 mg/ml respectively. The optimal

values from the GA were then inserted as inputs to the NN to yield an optimal ethanol concentration of 63.49 mg/ml.

The reliability of the ANN model was tested by comparing the predicted and experimental results. The observed ethanol yield from the experiments and the predicted ANN value were 64.71 mg/l and 63.49 mg/l respectively. In this case, there was satisfactory correlation between the predicted values and observed values from the experiment. The experimental analyses gave standard deviation of $\pm 1.1\%$.

5.3.2 Performance of the Genetic Algorithm

As explained earlier, once the neural network had been trained, it can be used to forecast ethanol production changes with changes in either temperature, sugar or yeast concentration without performing physical experiments. Thus, any multivariable optimisation method could be employed to seek the optimum ethanol product yield from the fermentation process. The genetic algorithm was used as the method of choice because it has a high convergence speed and its global optimisation ability (Zhu *et al.*, 2004; Ranganath *et al.*, 2002). Like the ANN, the genetic algorithm is a tool based on artificial intelligence. It aims to imitate the biological process of natural selection by cross-over and mutation (Holland, 1992; Krishnakumar and Goldberg, 1992)

Operation of the Genetic Algorithm

The theory behind the functioning of the genetic algorithm can be summarised as follows:

- An initial population is created. Since the variables have limits, the initial population must lie within these constraints.
- A fitness value is assigned to each member of the population in order to “score” it. The fitness function is used. Since the ANN model does not produce a function, but rather is trained to emulate the RSM empirical model, this equation (i.e. Equation 2.15) was used in the genetic algorithm as the fitness function and the returned values were then employed in the ANN model.
- Based on these fitness values, some members of the population are upgraded to become “parents” in the generation. Of the remaining members of the population, those with fitness scores that are relatively lower, are removed and placed into another population.
- Now the “parents” are made to produce children. As is the case in real-life genetics, two parents make one child. In the case of genetic algorithms however, this is performed by either mutating a single parent or by combining the vectors of a pair of parents-this is called crossover and is what has been used in this experiment.
- All the parents in the population are replaced in this way, by “new born” children who form the next generation.

- The cycle repeats itself (Goldberg *et al.*, 1989, Shimizu *et al.*, 1997), and it may be seen that newer generations will be more “fit” than the previous ones because the parents are selected based on their fitness values. The population will also decrease in number because the replacement of two parents by 1 child (formed during cross-over).
- However, the tool is not run until but one member is left in the generation, as this might take some time, it is run until one of seven stopping criteria has been met. These are: Generations, Time limit, Fitness limit, stall generations limit, stall time limit, and function tolerance level and nonlinear constraint tolerance.

Figure 5.10 gives a graphical illustration of the stopping criteria levels and which stopping criteria (of the above mentioned seven) caused the GA algorithm to stop running.

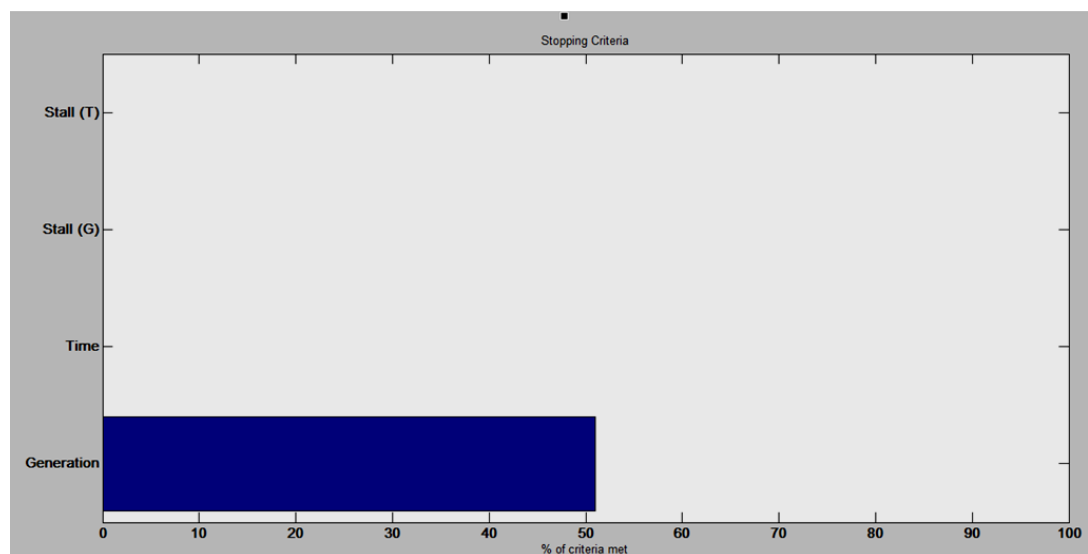


Figure 5.10: Levels reached by the stopping criteria

From Figure 5.10, it may be seen that the 50 iterations (generations) had been performed before the GA stopped running. It must be noted the maximum number of iterations/generations was set as 100.

The Rank scaling function was chosen and used because it “ranks” the fitness’ of individual. For example, the member with the highest fitness score is ranked 1; the next fittest member is ranked 2 and so forth. In this way, it removes the effect of the spread of these fitness scores and creating a hierarchy within the generation. In this way, it is superior to the “proportional” scaling function which has weaknesses associates with it and the “top” and “shift linear” scales functions, hence it was the scaling function of choice here.

As previously mentioned, the scaling function is used to convert the fitness scores into values that lie in the range of the selection function. The selection function is the means by which “fit” members are selected to move to the next generation to become parents and produce “fitter” children (by means of crossover). The stochastic function was selected because it has many advantages over the other selection functions. For instance, the “uniform” selection function which selects parents at random, this result in a somewhat undirected search.

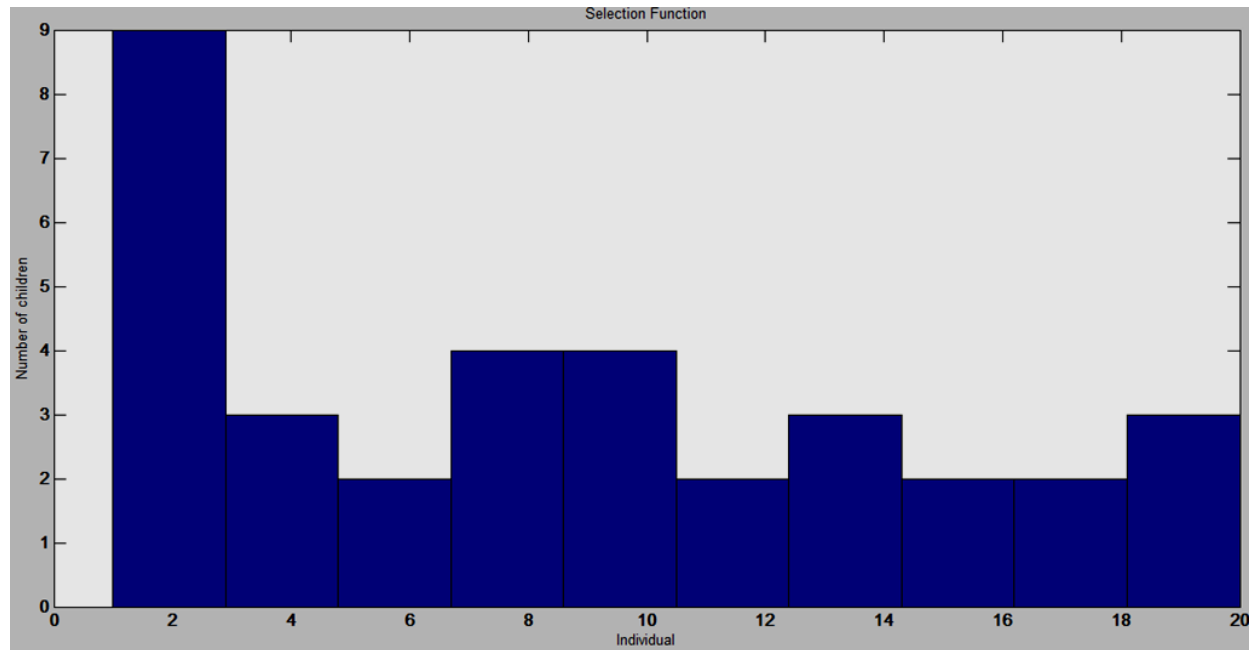


Figure 5.11: Histogram of the parents

Figure 5.11 shows the parents which are contributing to each generation, once the selection function has been applied. Once the selection function was specified, the reproduction procedure was set as follows: the minimum number of individuals that will be passed to the next generation (the “elite”) is set as two. This is because, in the crossover function of creating children, two parents will be needed, and the cross over fraction was 0.8.

For migration settings, *i.e.* when the individuals are passed over to the next generation, a forward scheme was selected (that is they should move in the forward direction of the algorithm).

5.3.3 Genetic Algorithm aspects during optimisation

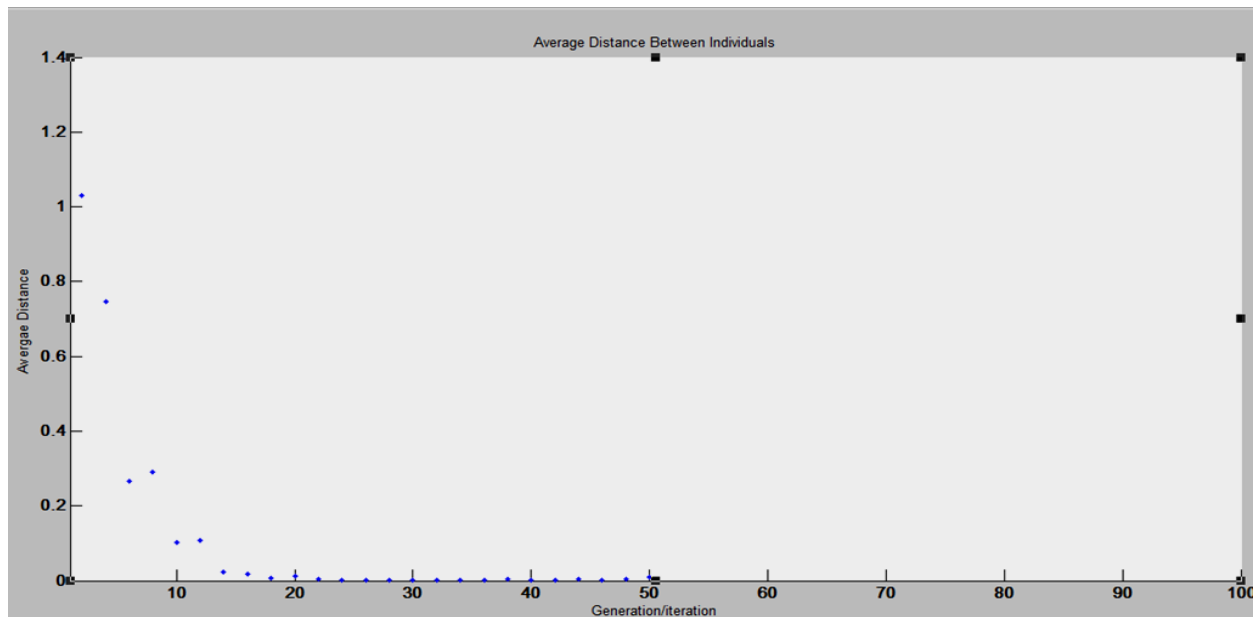


Figure 5.12: Average distance between individuals at each generation

Figure 5.12 asserts the results because as can be seen from the diagram, the distance between individuals was minimised from the 22nd generation (which is the 22nd iteration) to the 50th. This is analogous to a chemical process that is said to have reached its steady state because no further changes in the system (in this case the successful generations) are observed.

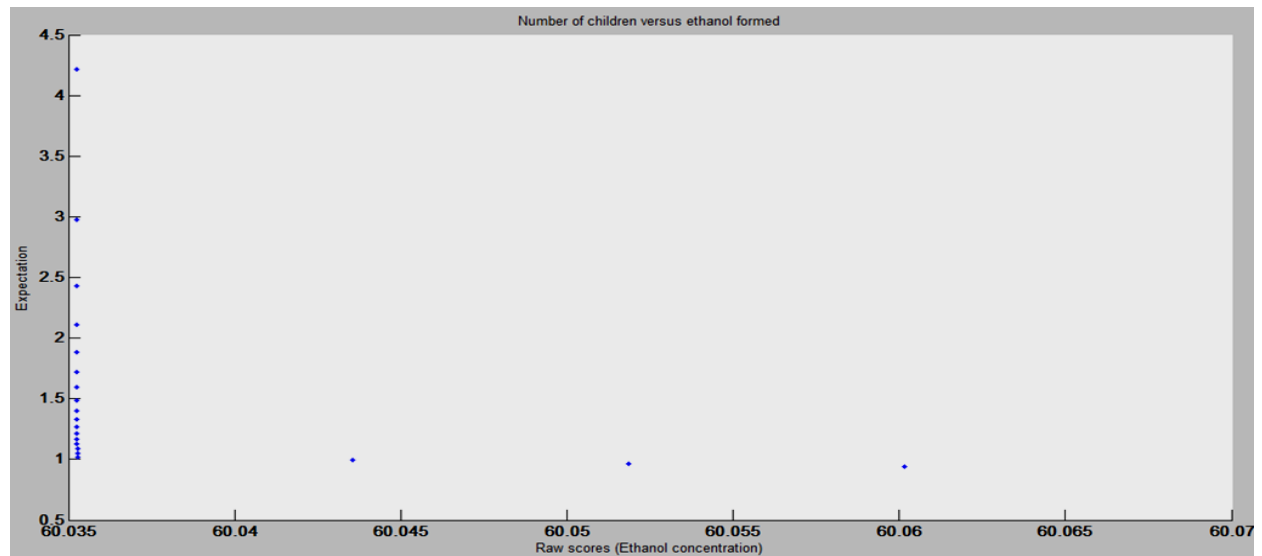


Figure 5.13: Expected number of children vs. raw scores at each generation

Figure 5.13 is a plot of the expected number of children versus the raw scores. Recalling that the raw scores are the fitness values of the individuals before fitness scaling for the selection purposes, this figure represents a good result. The fitness function gives the ethanol concentration as derived from response surface methodology. From the Figure 5.13, it can be deduced that from an ethanol concentration of (y-value of) approximately 60.04 mg/ml, the expected number of children is constant at 1. This means the selection function stops here because two individuals are necessary for the crossover function to move into the next generation to form another child (mentioned above as the “elite”).

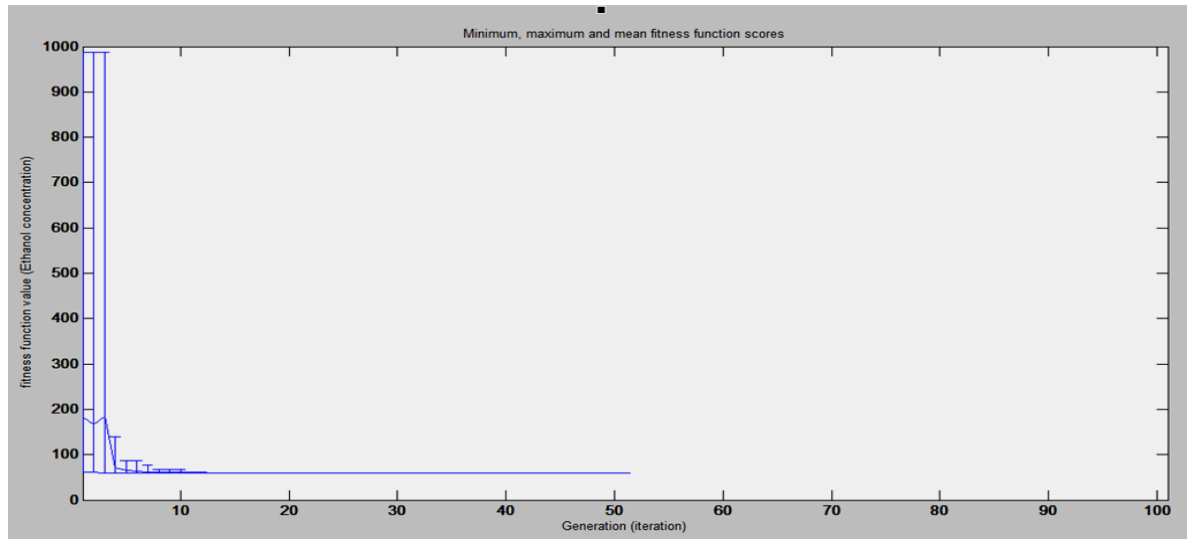


Figure 5.14: Best, Worst and Mean fitness function value (i.e. ethanol concentration) in each generation

Figure 5.14 is a plot of the minimum, maximum and mean (average) fitness function value. The fitness function was taken as Equation 2.15 because the trained neural network does not produce a fitness function but emulates Equation 2.15 (obtained from RSM) quite precisely. This graph confirms that the optimised ethanol concentration is 60.04 mg/ml because the mean fitness function value (i.e. the ethanol concentration) stabilises at 60.04 mg/ml from the 13th generation or so. Appendix D shows plots of scores which support the above mentioned results, which are the ANN performance graphs.

5.6 Supporting GA plots for an optimum Ethanol concentration

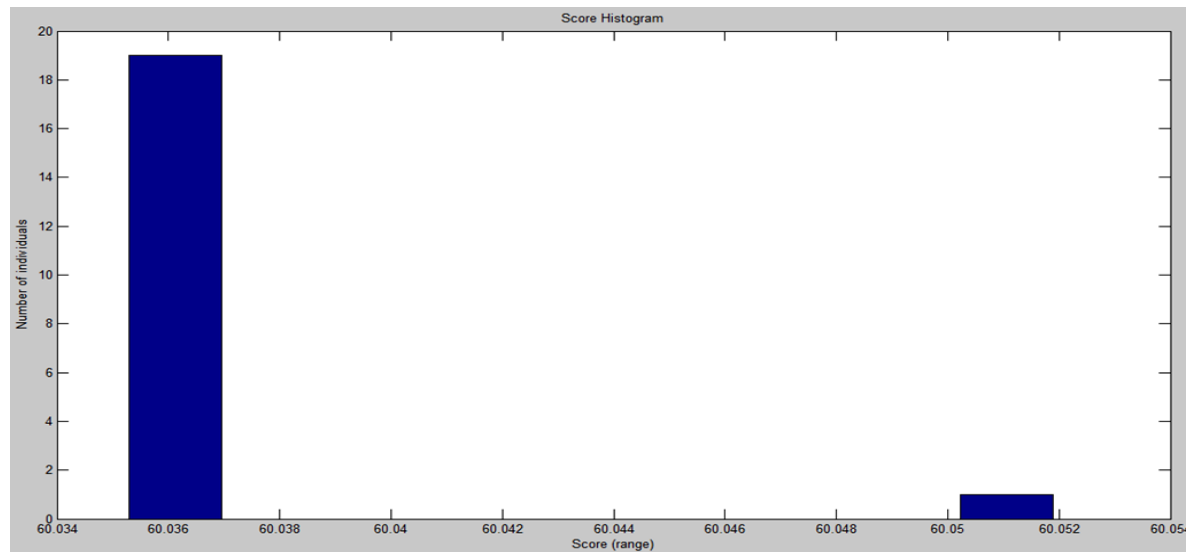


Figure 5.15: Histogram of the scores at each generation

Figure 5.15 gives an indication of the scores of the individuals at each generation. The number of individuals with the same fitness score are grouped together to produce the histogram seen above. Figure 5.15 can thus be seen as an alternative way of presenting the results concluded from Figure 5.13 because the number of individuals at which the score is 60.05-60.06 is still 1. That is, no more individuals can be taken to the next generation to reproduce and hence this may be taken as the optimum fitness function value, i.e. the optimum ethanol produced.

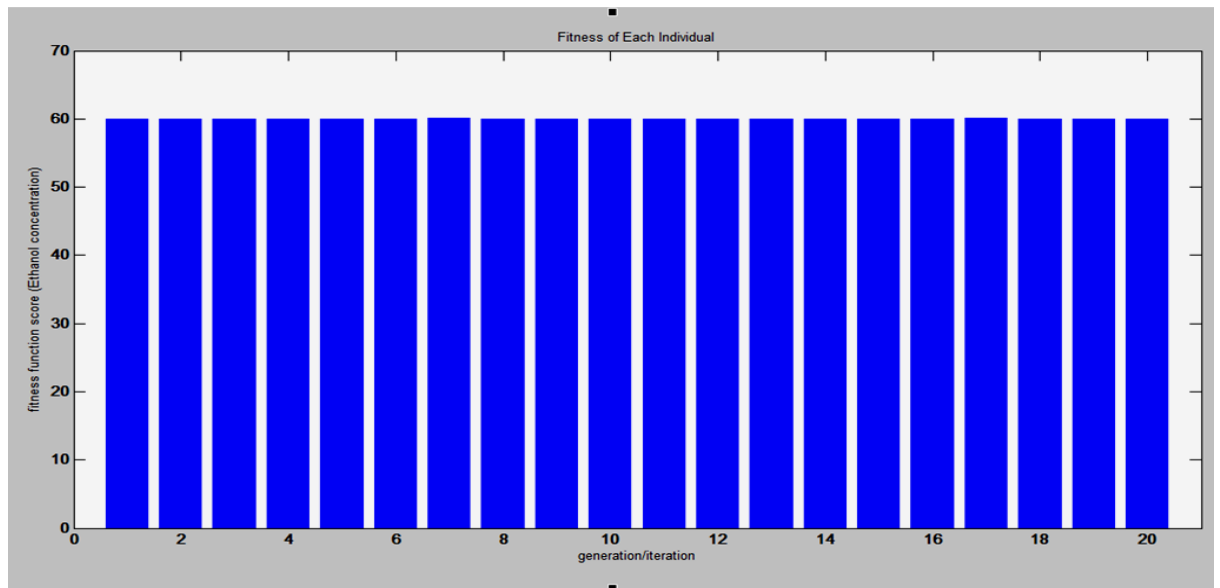


Figure 5.16: Scores of individuals at each iteration

Figure 5.16 shows that the fitness function value (i.e. y in Equation 2.15) is 60mg/ml for each individual which asserts the above mentioned deductions.

5.4 Summary and conclusions

The main objective of the study in this chapter was to model the fermentation process using Artificial Neural Networks (ANN), which could be used to optimise the ethanol yield during brewing. Data from the brewery was supplemented by data generated with RSM model as it requires a lot of data to model ANN. The results indicated the following:-

- Data from the RSM model is reliable to supplement experimental data in development of the ANN model.

- The ANN model proved superior and more accurate than RSM model when tested against various experimental results.
- Architecture of the ANN model consisted of three layer (3 input: 10 hidden layer: 1 output) feed-forward back-propagation neural network with Levenberg-Marquadt method.
- Use of ANN with Genetic Algorithm (GA) function optimiser yielded satisfactory and accurate results.
- The adequacy of the developed ANN model proved satisfactory when validation checks were performed.

CHAPTER SIX

Integration and operation of the Wits Micro-brewery

“To improve is to change; to be perfect is to change often ”

- **Winston Churchill**

6. INTEGRATION AND OPERATION OF WITS MICRO-BREWERY

6.1 Comparison of batch and fed-batch fermentation

6.1.1 Introduction

Improved operational methods can yield optimum process results. Most breweries still use batch fermentation process mainly because of their reluctance to change and lack of sufficient data illustrating benefits from other alternatives. Challenges in the brewing industry such as competitive product for the desired cost reduction are forcing industry to explore new production options.

More efficient production strategies of beer are required to reduce input waste and maximise production outputs. Fed-batch mode of operation has been successfully used for production of pharmaceuticals, and other biotechnological products required in large quantities (Nilsson *et al.*, 2001). This section is about the comparison of batch and fed-batch fermentation based on ethanol yield. As noted from the RSM model in chapter 4, temperature is the significant influential factor during fermentation; hence the study was based on their performances at varying temperatures.

6.1.2 Materials and Methods.

The main aim of the investigation was to compare the batch and the fed-batch modes of fermentation i.e. investigate the feasibility of the fed-batch mode relative to the traditional (batch) mode of fermentation at Wits Microbrewery plant. The first set of results draws a comparison between the two at different operating temperatures; the second set compares the batch and fed-batch fermentation findings to the theoretically obtained values to check if the estimated feeding point is optimal for the process.

6.1.3 Yield and productivity of batch and fed-batch fermentation at various temperatures.

The following results show yield and production rates obtained from batch and fed-batch fermentations at various temperatures. The pH for this part of experiment was kept at optimum pH of 5 as was reported in our earlier investigation (Madigoe, 2009). Figure 6.1 shows results obtained at 10°C.

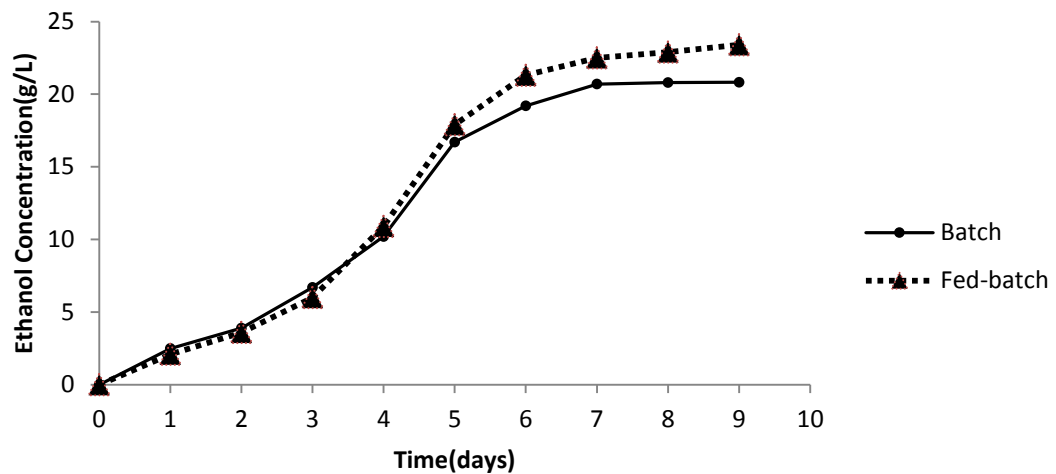


Figure 6.1: Ethanol productivity for batch and fed-batch fermentation at 10°C

From Figure 6.1, ethanol production rate (g/l.hr) at 10°C is the same for batch and fed-batch process during the first four days. This is despite the fact that the batch fermentation has twice the volume of the substrate than the fed-batch fermentation within that period. The reason for this might be the fact that before aggressive cells multiplication or fermentation starts, the cells need to adapt to the environment first (Iyuke *et al.*, 2008), mostly because 10°C is below the minimum optimal performing temperature for *Saccharomyces Cerevisiae* (Cinar, 2003). Another reason for this behaviour might be the impression that the *Saccharomyces Cerevisiae* cells are overloaded. This is in agreement with Jong *et al.* (1995), who stated that when high concentrations of substrate are present, the cells get "overloaded", that is, the oxidative capacity of the cells is exceeded, and due to the Crabtree effect, products other than the one of interest are produced, reducing the efficacy of the carbon flux. However, after the fourth day the fed-batch process shows a higher ethanol production rate of 0.104 g/l.hr when

compared with that of batch process at 0.088 g/l.hr. Although the cells were “overloaded” initially, during fermentation process they multiply in numbers. The phenomena observed after day four indicates that the cells are “underfed” for batch fermentation mode. Addition of substrate (wort), during the fourth day into the fed-batch provided sufficient sugar to be catalysed into ethanol, also increased contact surface area between *Saccharomyces Cerevisiae* and the wort hence resulting in higher fermentation rates. Fed-batch fermentation ended with higher ethanol yield of 23.4 g/l than the final yield for batch fermentation of 20.82 g/l of ethanol. The poor performance of batch fermentation can be attributed to the fact that both underfeeding and/ or overfeeding of the cells can be detrimental to cell growth or product formation (Dombek and Ingram, 1987, Lee *et al.*, 1999). Figure 6.1 indicates that for both modes fermentation was completed at day 8.

Results for fermentation at 20°C are shown in Figure 6.2 below.

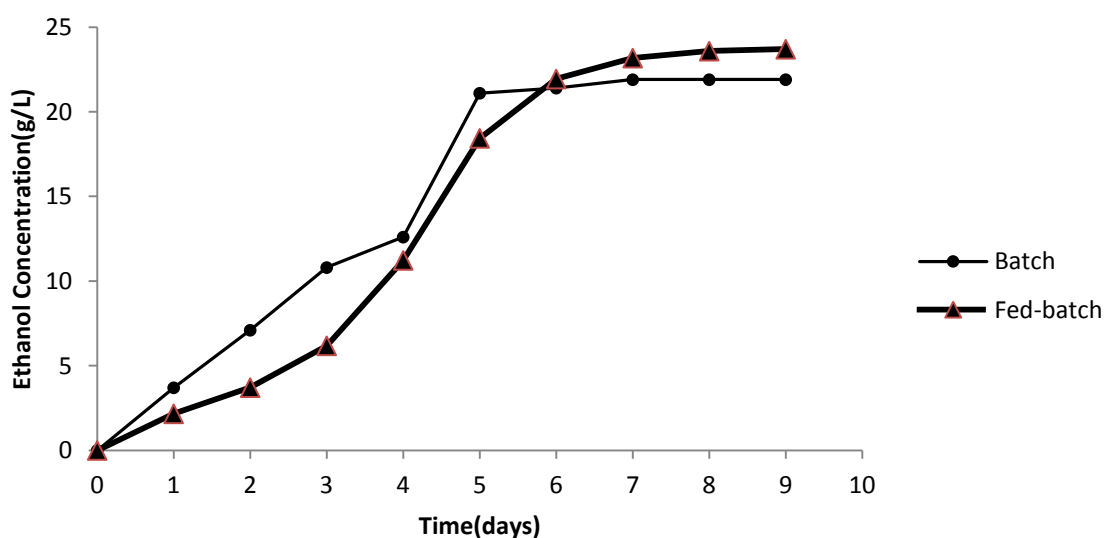


Figure 6.2 Ethanol productivity for batch and fed-batch fermentation at 20°C

At 20°C, similar trends as those observed at 10°C are visible. However, notably, batch fermentation starts with a higher production rate of 0.130 g/l.hr as compared to initial production rate for fed-batch mode which is at 0.104 g/l.hr. This indicates that the cells did not take long to adapt to the environment because 20°C is favourable fermentation temperature for *Saccharomyces Cerevisiae* (Manzano *et al.*, 2005). Due to favourable conditions, batch mode cells quickly started to multiply and because of available wort they catalyse more sugars into ethanol than fed-batch fermentation. Batch fermentation completes at day five whereas fed-batch fermentation continues until day eight. Early termination of the batch fermentation is due to high concentration of alcohol which can inhibit cell growth, as noted by Lee (1999). This highlighted the advantage with fed-batch, in that the feed of the required components for growth and/or other substrates required for the production of the product can never be depleted and the nutritional environment can be maintained approximately constant during the course of fermentation. Fed-batch advantage was also reported by Jong *et al.* (1995), whereby, during their investigation they noted that the production of by-products that are generally related to the presence of high concentrations of substrate can also be avoided by limiting its quantity to the amounts that are required solely for the production of the biochemical.

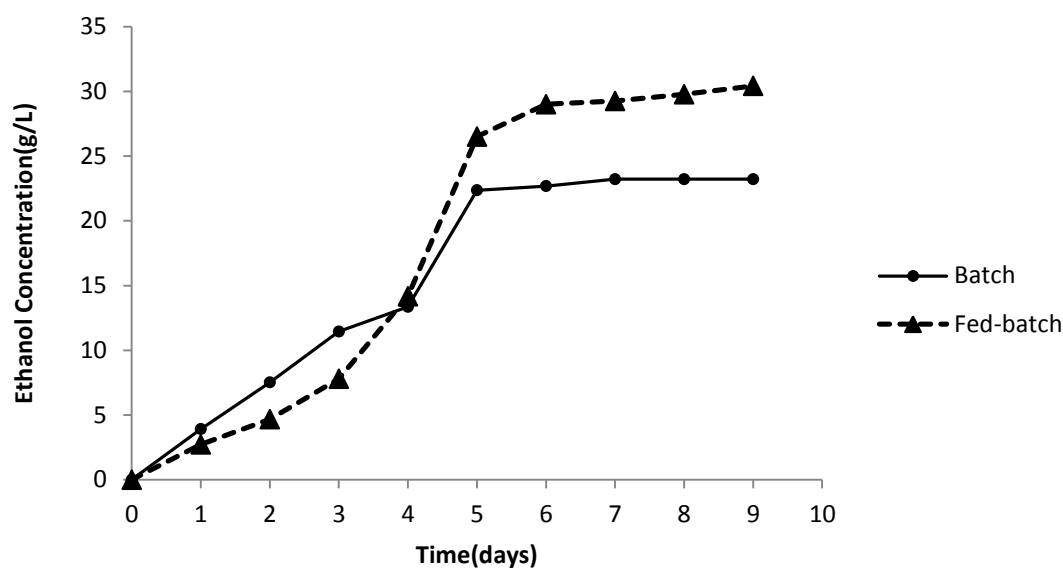


Figure 6.3: Ethanol productivity for batch and fed-batch fermentation at 30°C

Figure 6.3, where fermentation occurs at 30°C, further shows that the disadvantage of batch fermentation is that the rate of ethanol production per milligram of cell protein is maximal for a brief period, early in the process but declines progressively as ethanol accumulates in the surrounding broth. This is in agreement with what was observed by other researchers (Dombek *et al.*, 1987). The introduction of the second volume of wort for fed-batch during batch exponential growth induced an efficient enzyme production and resulted in higher yields of ethanol for fed-batch. Fed-batch ended with a significantly higher ethanol yield of 30.42 g/l, which is significantly higher than 23.21 g/l ethanol from batch fermentation. Figure 6.4 shows results for fermentation at 35°C.

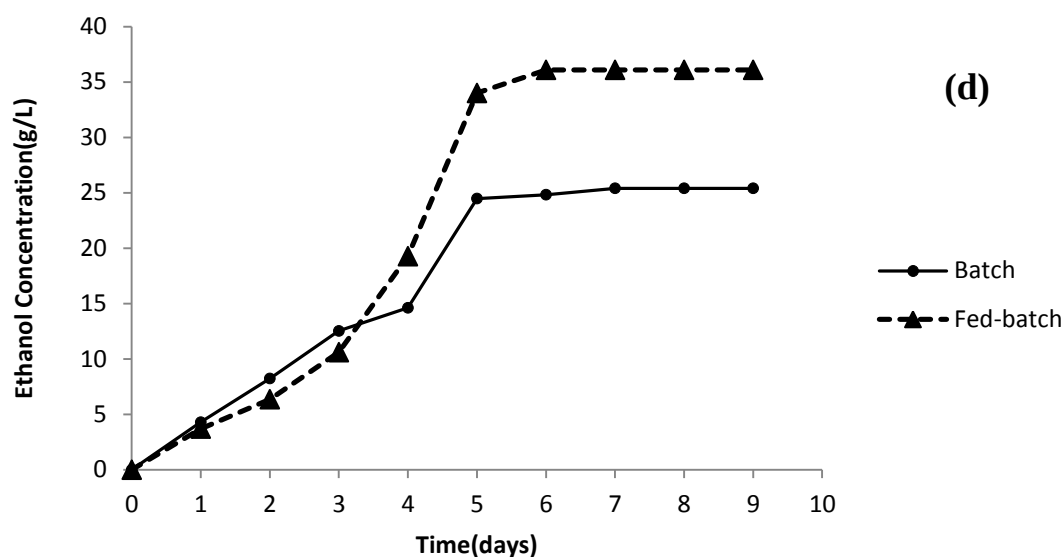


Figure 6.4: Ethanol productivity for batch and fed-batch at 35°C

At 35°C, the performance of fed-batch is superior to that of batch fermentation after 3.5 days of fermentation period. Fed-batch has an average production rate of 0.270 g/l.hr ethanol as compared to 0.150 g/l.hr ethanol from batch fermentation. There is also a highly significant difference in their yields, batch fermentation ended with low yield of 25.04 g/l as compared to high yield of 36.09 g/l ethanol from fed-batch fermentation.

The results indicate that as the temperature increases, the production rate also increases. The trend displayed by increase in fermentation rates with temperature in agreement with the observed experimental results reported in literature (Priest and Stewart, 2006; Phisalaphong *et al.*, 2006; Zhong *et al.*, 1995). The reason for an increase in ethanol yield due to temperature is the same as that discussed above

were by increasing temperature results in an increase in fermentation rate by increasing the rate of yeast metabolism giving higher specific fermentation rates.

6.1.4 Comparison between experimental and theoretical data correlations.

The optimal design of fed-batch cultivation has to take several factors into consideration, for example, the time to start the fresh nutrient feed (when the culture is exponentially growing) what should be the substrate concentration in the feed and its rate of addition and when to finish the nutrient feeding so that the highest concentration of product is produced and that there is minimal unconverted substrate when the reaction is complete. At this stage it is very important to check the quality of beer in terms of maturity, which is very important factor for the final consumer. To achieve this, a batch mathematical model was therefore developed using the kinetics of the fermentation which were then extrapolated to fed-batch cultivation by incorporating the mass balance terms in the model equations as shown in Chapter 4. To validate this assumption, experimental data was correlated to theoretically predicted data from the models in Chapter 4. Table 6-1 presents standard deviations between experimental data and theoretical data for batch fermentation and fed-batch fermentation.

Table 6-1: Standard deviations between experimental and theoretical data for batch and fed-batch fermentations

Temp (°C)	Batch STD (%)	Fed-batch STD (%)
10	19.40	24.70
20	14.30	22.10
30	9.80	14.20
35	9.30	14.10

When the experimentally obtained data is compared to theoretically predicted data, it can be observed that the two correlate reasonably well for non-linear biochemical process such as brewing fermentation. As Shimizu (1993) pointed out, model development for fermentation is not straightforward due to non-linear nature of bioprocesses and lack of accurate models describing cell growth and product formation, together with slow process response. It can also be noted that higher deviation are visible for fed-batch experimental and theoretical data. This can be attributed to the assumptions made when developing the model for fed-batch fermentation (i.e. that the production rate data available from batch fermentation can be applied to fed-batch (Longobardi, 1994). Depletion of the fermentable sugars and/or reduction of cells during the death phase can be another contributing factor to relatively high standard deviation. High deviation at low temperature for batch fermentation and fed-batch fermentation can be due to unfavourable environment for fermentation.

6.1.5 Summary and conclusions

The main objective of comparing batch and fed-batch fermentation was to evaluate the prospect of switching to fed-batch fermentation in the brewing industry in order to optimise ethanol yield during production. As temperature is the most significant variable for fermentation, the evaluation of the processes was studied at various temperatures. The following were observed during the study:-

- Fed-batch fermentation produced higher ethanol yields than batch, although the differences were not surprisingly major to justify the switch to fed-batch fermentation.
- Fed-batch fermentations had longer fermentation times than batch at all temperatures. This could lead to increased production length when implementing fed-batch operations.
- Successful investigation from the study indicated that fed-batch could be implemented in the micro-breweries without additional equipment needs.
- Increased complication of the fed-batch might be brought by possible spoilage of wort waiting to be added it were to be implemented industrially.

6.2 Integration of the Wits Micro-brewery plant

6.2.1 Introduction

Batch processes lack efficiency in energy optimisation as opposed to continuous processes and an increase in energy cost over the past few years made it necessary to save energy whenever possible (Majozi, 2006). Heat integration is the synthesis of heat exchanger networks to connect process streams that require cooling to those that require heating, with the primary objective of optimizing energy usage. Although heat integration has been successfully applied to continuous processes, it has posed challenges when applied to batch processes. This is due to the fact that heat recovery is not only limited by the thermodynamic properties of the streams such as temperatures and heat capacity flow rates, but also by time (Kemp *et al.*, 1988). To mitigate the limitations various methods that divide process streams in to both, temperature and time intervals has been proposed (Bozan *et al.*, 2001).

This section will contribute to present status of knowledge in batch process integration. A systematic approach is presented in this section that involves scheduling of batch brewing process operated cyclically, involving policy of no-intermediate storage.

6.2.2 Materials and methods

6.2.2.1 Integration using rescheduling method

In batch processes, many streams are available for a particular period of time, which restrict the possibility of heat exchange. Also, for most batch systems, much heating is done *in-situ*; this makes it impossible for them to run at constant temperatures with constant heat capacity flow rates. In order to incorporate the time constraints in batch process integration, a scheduling approach has been applied. With this approach, streams from one time interval are transferred to another time interval(s). This method of process integration is termed rescheduling method. The approach is explained below using the Wits Microbrewery as case study.

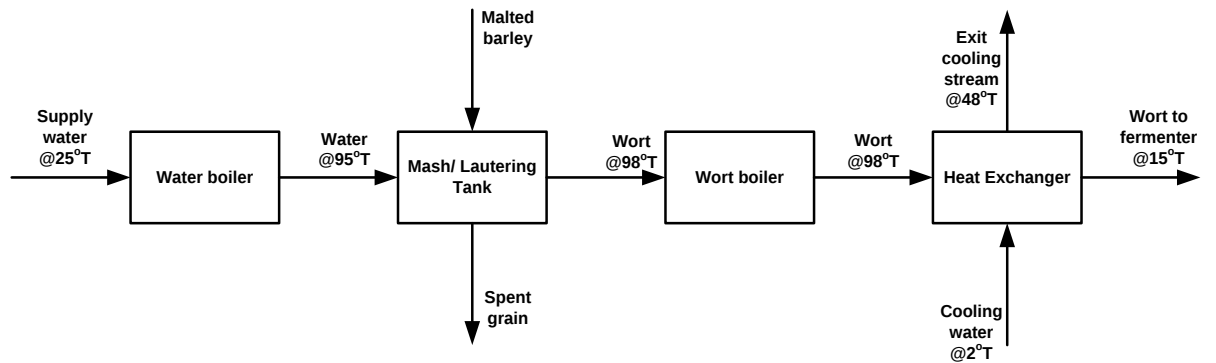


Figure 6.5: Simplified process flow diagram for batch brewing at Wits Micro-brewery plan

Figure 6.5 shows the flow diagram of a single product (beer) plant comprising of boilers and a heat exchanger. The fermentation tank/ process is excluded as it runs

after the production for a period of time. The process operates in the following way: raw material water (60 L) is available at 25°C and fed to the first boiler where it is heated to 95°C using electrical coil. The boiled water is discharged to the mash/ whirl pool tank for mashing. In the whirlpool/ mash tank, the boiled water is combined with feed, malted barley, for mashing. The mixture of boiled water and malted barley is left in the mash tank for 75 minutes to dissolve sugars from the malted barley. The exit temperature of dissolved sugars in water, now called the wort, is transferred in to the second boiler. Wort supplied at 65°C is now boiled to 98°C for a period of 1.5 hours. Thereafter, the boiling wort is given a heat shock via the heat exchanger before being discharged into the fermentation tank. Wort exits the heat exchanger at 15°C, which is desirable for the required fermentation. Cooling water for the wort is supplied at 2°C from the separate refrigeration system and exit the heat exchanger at 48°C into the drain. The plant operating schedule is given in Table 6-2. The streams are shown in Table 6-3 where they are classified as hot or cold. The streams requiring cooling are termed “hot” and the streams requiring heating are termed “cold”.

Table 6-2: Plant operating schedule

Stream no.	Supply T°C	Target T°C	Heat Flow (kJ/min)	Start time (min)	Finish Time (min)	Required/loss utility (kJ)	Stream description
1	25	95	294	0	60	17640	Feed water boiling
2	95	65	-100.8	60	135	-7560	Wort Cooling during mashing
3	65	98	92.4	135	225	8316	Wort boiling for hop addition
4	98	15	-2091.6	225	235	-20916	Cooled Wort to fermenter

Table 6-3: Stream data for the case study

Stream no.	Type	Supply T°C	Target T°C	Heat Flow(kJ/min)	Start time(min)	Finish Time(min)	Required/loss utility(KJ)
1	Cold A1	25	95	294	0	60	17640
2	Hot B1	95	65	-100.8	60	135	-7560
4	Cold A2	65	98	92.4	135	225	8316
6	Hot B2	98	15	-2091.6	225	235	-20916

The volume of reactants and products is 60 Litres and the specific heat capacity is 4.2 kJ/kg.K. The maximum achievable cooling is dependent upon the minimum approach temperature used as well as the size of the heat exchanger, for practical reason an acceptable ΔT_{min} of 10°C is used (Coulson *et al.*, 1999). For batch processes, total heat transferred (J) to a liquid in the vessel, Q, to the liquid is governed by;

$$(6.1)$$

where

m_L = mass of liquid in the vessel (Kg); C_{pL} = specific heat capacity of liquid (J/kg.K); T_{L2} = final temperature of the liquid in the vessel (K); T_{L1} = initial temperature of the liquid in the vessel (K).

From this it can be seen that the existing process, with no heat integration, requires 25956 KJ of hot utility and loses 28476 KJ of potential hot utility. As it

stands, time events of the process streams can be represented as shown in Figure 6.5.

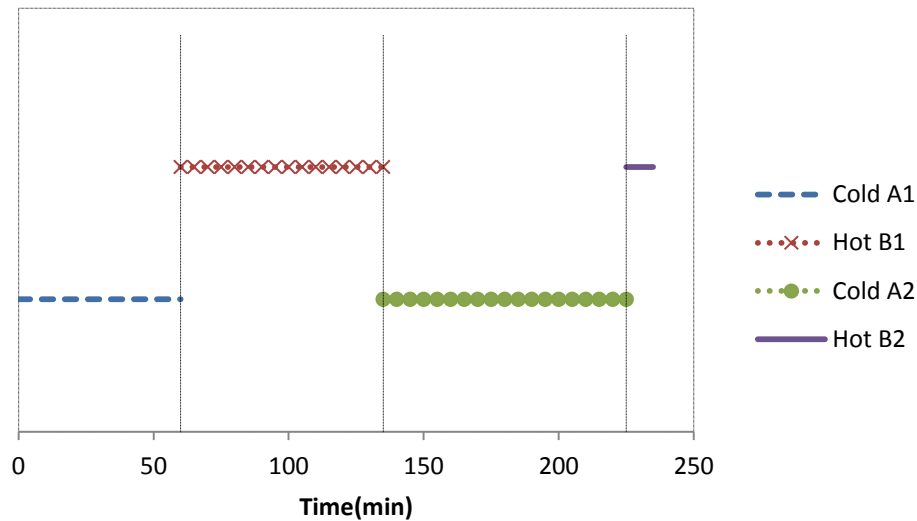


Figure 6.5: Time event charts for the streams

From the time event chart, it is clear that those hot and cold utilities are never required simultaneously. Clearly with the process as it stands, no heat recovery through integration is possible. However, when another production is required and hence the process needs to be repeated, heat may be recovered into the next batch. In this case study, it was assumed that two productions were required on the same plant. Thus process integration was performed for the two productions. However, the same principle could be extended for any number of productions.

6.2.3 Result and Discussion

The case study that was considered for batch process integration was production of beer in the Wits Micro-brewery using optimised conditions from the ANN and

RSM Models together with optimised mashing process from Iyuke *et al.* (2008). Although the production times and energy savings might vary depends on the specific brew, they represent the major characteristics of the typical batch production process in the Wits Micro-brewery plant.

As explained, from the Wits micro-brewery plant, hot and cold utilities are never required simultaneously; and this is the case with most batch processes. However, when another production is required and hence the process needs to be repeated, heat may be recovered into the next batch. In this case study, it was assumed that two productions were required on the same plant for production of the same product. Thus process integration was performed for the two productions. However, the same principle could be extended for any number of productions. Process integration, using rescheduling technique was used in the Wits micro-brewery plant. The suitability of the rescheduling method was demonstrated by applying it in food industries such as the cheese industry, the paper mill, iron and steel, textiles and cement (Becker *et al.*, 2002). From the scheduling data in Table 6-1, it was found that four time intervals exist in the process, each with its stream population. Figure 6.6 shows the temperature profile of each stream with time for the data represented in Table 6-2.

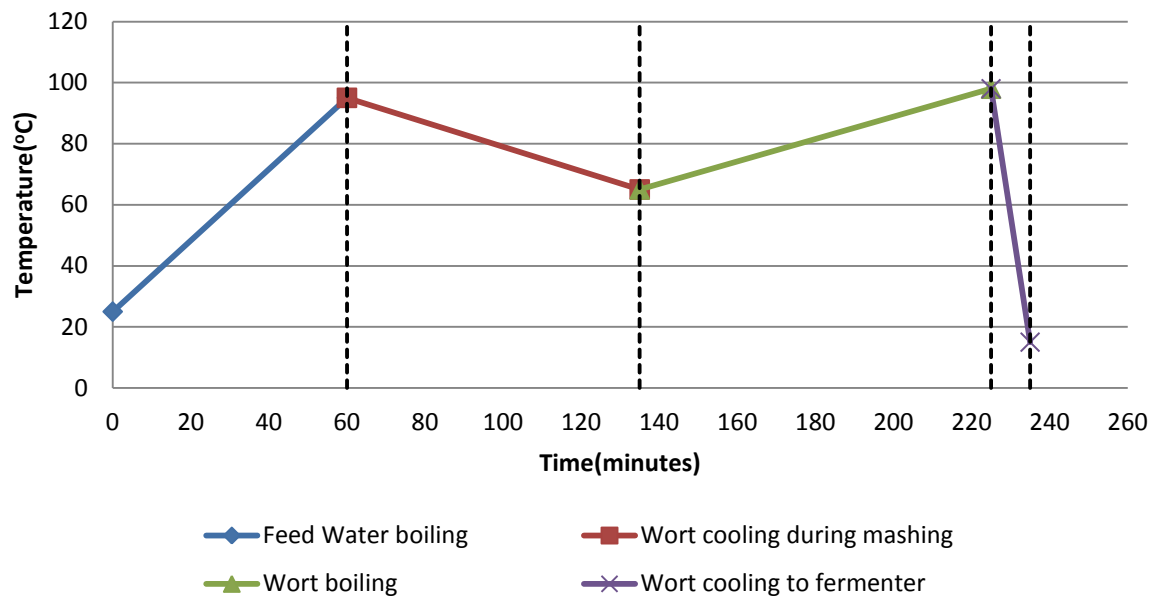


Figure 6. 6:Temperature profiles for batch process

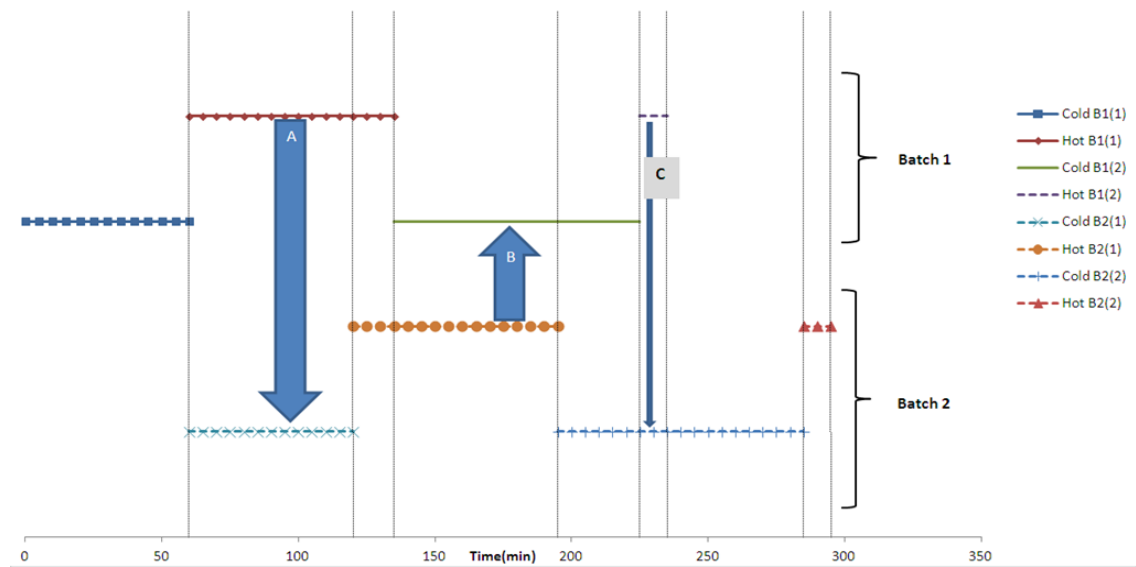


Figure 6.7: Time event chart for rescheduled process

Figure 6.7 shows the streams of the first brew and an identical second batch, which could be the batch after the first one, using the same equipment. If the second batch can start after 60 minutes, its initial cold stream could be integrated with the hot stream from the first batch. As indicated earlier, process integration through changing the process timing in this manner is called rescheduling. If instantaneous heat flows of the streams are mismatched, load limitations may affect the optimum heat that can be recovered after integration of streams. Instantaneous heat flows for HotB1 (1) and Cold B2 (1) are -100.8 KJ/min and 294 kJ/min. That means that hot stream from the first batch can only supply 100.8 kJ/min of heat flow and thus an additional 193.2 kJ/min is required to heat the Cold B2(1) to its target temperature.

Time slots 120-225 min shows that Cold B1 (2) requires 92.4 kJ/min while Hot B2 (1) release 100.8 kJ/min of energy; which means the remaining 8.4 kJ/min of energy goes to waste. Similar analysis can be performed for all the time slots. Multiplying the heat net heat loads by the lengths of time periods, after process integration by rescheduling, energy requirements for the consecutive process can be calculated. For this case study, the following are the findings:

(a) Heat integration calculations for Integration A

Energy recoverable from stream Hot B1 (2) = $-100.8 \text{ kJ/min} \times 60 \text{ min} = 6048 \text{ kJ}$, and because of load mismatch, an addition of $(294-100.8) \text{ kJ/min} \times (120-60) \text{ min} = 11592 \text{ kJ}$ is required.

(b) Heat integration calculations for Integration B

Heat loss for time slot 120-135 minutes (stream Hot B2 (1)) is $100.8 \text{ kJ/min} \times (135 - 120) \text{ min} = 1512 \text{ kJ}$. For time slot 135-195 min, cold stream, Cold B1 (2) requires 92.4 kJ/min ; which means energy recoverable is $92.4 \text{ kJ/min} \times (195 - 135) \text{ min} = 5544 \text{ kJ}$.

It is clear that heat loss from Hot B2 (1) is $(100.8 - 92.4) \text{ kJ/min} = 8.4 \text{ kJ/min}$; which translate into $8.4 \text{ kJ/min} \times (195 - 135) \text{ min} = 504 \text{ kJ}$.

(c) Heat integration calculations for Integration C

For time slots 195-225min, Cold B2 (2) requires $92.4 \text{ kJ/min} \times (225 - 195) \text{ min} = 2772 \text{ kJ}$. Heat required for Cold B2 (2) for time slot 225-235 min is $92.4 \text{ kJ/min} \times (235 - 225) \text{ min} = 924 \text{ kJ}$. An additional heat of $92.4 \text{ kJ/min} \times (285 - 235) \text{ min} = 4620 \text{ kJ}$ is required for time slot 235-285 min.

Heat loss/ waste of $2091.6 \text{ kJ/min} \times (295 - 285) \text{ min} = 20916 \text{ kJ}$ is experienced for time slot 285-295 min. This means that after integration, energy required for the second consecutive batch is 19908 kJ and heat loss is 22932 kJ . Energy loss from Stream Hot B2 (2) could be further recovered with the next batch to follow.

It should be noted that, the potential heat recovery/ exchange mention above is achievable if heating and cooling are not done in the same vessel. Hot liquid needs to be discharged to a separate vessel or heat exchange so that it can heat up the next batch. From the heat integration targets performed, it can be noted that

heat recovery can vary between 0-100percent, with a number of intermediate options.

Other option such as heat storage can be deployed to recover heat losses. However, heat storage in practice is rare and likely to remain so. This can be due to the fact that it is difficult to store heat for the relatively long periods that batch processes undergo. They result in unacceptable heat losses; and most storage methods require much larger driving force than ΔT_{min} , thus, they tend to be more costly than heat exchange. First batch without process integration had energy loss and requirement of 28476 kJ and 25956 kJ, respectively. After heat integration, heat loss and requirement were reduced to 22932 kJ and 19908 kJ, respectively. That means, after heat integration process was incorporated to the Wits Microbrewery plant, heat loss and energy demand were reduced by 19.5% and 23.0%, respectively.

6.2.4 Summary and conclusions

Energy saving, global warming and greenhouse gas emissions have become critical socio-economic issues that scientists have to solve using innovative solutions (Linnhoff, 1982). Process integration has the potential to play a pioneering role in solving the related problems. The main objective of heat integration in the study was to optimise the brewing production by improving the existing situation through reduction in energy demands and waste. It consists of matching hot streams with cold streams within the process. Challenge in integration of batch processes is due to the fact that cold and hot streams are not always available at the same time, unlike in their counterparts continuous

processes. That means to improve energy situation in batch processes, time and temperature becomes two most important variables. Process rescheduling coupled with heat integration was successfully applied to brewing process operated cyclically and the following were observed:-

- Heat recovery in batch process integration is possible but with minimal savings;
- Besides time and temperature constraints, energy load in mismatch is one of the critical variables that restrict maximum energy recovery in batch process integration; and
- For maximum benefits during cyclical operations, it is important that plant operators work efficiently as batch process integration is sensitive to time and temperature loss.

CHAPTER SEVEN

Conclusion

“When you cease to make a contribution, you begin to die”

- Eleanor Roosevelt

7. CONCLUSIONS AND RECOMMENDATIONS

7.1 Conclusions

7.1.1 Introduction

The brewing industry is one of the widely practised processes, mainly due to its product demand and its economic contribution. Despite technological improvements over the last two decades, energy consumption and increasing costs of feeds has forced the scientists and engineers to come up with innovative solutions to make the industry competitive. Temperature, yeast and fermentable sugars from the malted barley are the main input variables that affect brewing fermentation. Optimisation of this process variables coupled with energy savings can go a long way in making the industry sustainable.

The main objective of this thesis was to optimise Wits Micro-brewery plant using Artificial Neural Networks (ANN) and integration of the process. A review of literature indicated that brewing process exhibits complex and non-linear behaviour which proves it difficult to model and optimise with basic linear models. Process optimisation models such as Response Surface Methodology (RSM) have proved successful in other instances as it takes into consideration interaction between process variables; which is ignored in ‘one-at-a-time’ optimisation methods. If RSM is proven to be closely reliable to ANN models, it can, in most cases be applied for optimisation of complex bio-process instead of

applying superior ANN models which requires a lot of data and could be heavily sophisticated for most engineers. Most breweries are reluctant to use fed-batch fermentation, although fed-batch fermentation has been successfully applied and yielded higher outputs in pharmaceuticals and other bio-industries. Thus, the prospect of switching to fed-batch fermentation in the brewing industry has the potential to optimise resources and reduce waste. Heat integration in batch brewing processes can play a pioneering role in reduction of energy demands and waste through heat recovery.

In order to investigate the possibility of process optimisation and integration, the aims of this study were defined as:-

- (i) Develop Artificial Neural Network (ANN) model that relate selected process variables to performance for complex brewing fermentation system;
- (ii) Investigate the reliability of RSM for analysis and optimisation of selected variables and compare it to ANN;
- (iii) Comparison of fed-batch fermentation vs. batch fermentation; and
- (iv) Integrate the brewing process.

7.1.2 Optimisation using Response Surface Methodology

The model relating medium concentration, temperature and yeast concentration during fermentation was successfully formulated using Response Surface

Methodology. It has been shown that the application of RSM for optimisation of brewing process can be an interesting alternative to one variable at a time optimisation. An additional advantage presented by RSM was the ability to study the interactions among all the three components investigated. It was found after analysing and having generated the fitted model that the model predicted an optimal ethanol production conditions at temperature of 21.7°C, sugar concentration of 164 mg/ml and yeast concentration of 0.35 g per 500ml of solution. The fitted model was tested and found to have a confidence level of 99% and statistically 82% of the sample variation could be explained by the model. Confirmatory results showed that the model fits the experimental data quite well, with an error margin of 12%. It was thus concluded that the fitted model predicts accurately the ethanol yield with an accuracy margin of 80%. The results critically pointed out the importance of temperature for brewing fermentation.

7.1.3 Optimisation using Artificial Neural Networks

Optimisation of fermentation process was investigated using Artificial Neural Network with the integration of Genetic Algorithm. A three layer Feed Forward Neural Network with ten neurons in hidden layer integrated with Genetic Algorithm was effective in predicting the optimal ethanol concentration from an experimental model developed using Response Surface Methodology. The Genetic Algorithm identified the optimal coded input values using the experimental model which were consequently fed as inputs to the Neural Network after training. The optimal ethanol yield was 63.49 mg/ml with corresponding

process input variables of 19.0°C, 163.9 mg/ml and 0.42 mg/ml for temperature, sugar and yeast concentration respectively. The ANN model was able to predict optimised ethanol yield with 98% accuracy. It was thus concluded that ANN is suitable for prediction and optimisation of complex and non-linear processes like brewing fermentation.

It is recommended that the limits of the temperature variable in the RSM be changed to be restricted to only between . Although this narrow range is only applicable for the yeast to be used with ale fermentation, it also reduces the range in which the variables are limited. This will enhance the GA in its global search for optimum variables.

7.1.4 Comparison of batch and fed-batch fermentation

Fed-batch fermentation during brewing process can produce higher yields and higher production rates as compared to batch fermentation. The results critically pointed out the importance of temperature for brewing fermentation as the production rates and the final yields for both batch and fed-batch increased with an increase in temperature. Experimental and theoretical data for batch fermentation and fed-batch fermentation correlated satisfactorily well. As shown from the results, it can then be concluded that, utilization of fed-batch fermentation as compared with batch fermentation implies increase in productivity. The next focus of our research is on optimisation of the fed-batch fermentation process.

7.1.5 Integration of the brewing plant

After heat integration was deployed to the Wits Microbrewery plant, heat loss and energy demand were reduced by 19.5% and 23.0%, respectively. It can thus be concluded that heat recovery in batch processes is possible, but generally gives lower savings when compared to continuous processes because of constraints that result when hot and cold streams do not exist at the same time. Opportunity for substantial heat recovery is also negatively affected by load mismatch between integrated streams. Rescheduling, as an integration technique for time-dependent and batch processes, has been demonstrated as a powerful tool in evaluating heat integration possibilities.

7.1.6 Potential application in industry

The study has presented new and efficient method for brew optimisation that integrated three most important and significant variables in the brewing process, namely, temperature, yeast concentration and carbohydrate concentration. Developed models will go a long way in helping brewers predict and optimise production yields. Results from the investigations will contribute to knowledge and relationships in brewing (and related bio-processes) with respect to ethanol yield, which is the product of interest, and important input variables as indicated. New knowledge and direction into possible application of fed-batch fermentation in the brewing industry has been paved. The research has also added knowledge

to the on-going studies on integration of batch processes, which are important for energy savings, waste management, and environmental sustainability.

7.2 Recommendations

The study has assisted in the understanding of optimisation of brewing process together with recommendations that need further investigation.

Optimisation using Response Surface Methodology and Artificial Neural Network

In this study, RSM and ANN optimisation techniques were applied to brewing process. It was shown that ANN possesses superior accuracy than RMS, in prediction, and thus optimisation of brewing fermentation. Although these methods were developed on beer production, the bigger picture is the development of the general framework which is transferable to other types of fermentations, and trials with other fermentation process are encouraged in the future. This is particularly important since the input variables that were looked at in this research are considered to the most significant factors in fermentation; and were only looked at simultaneously for the first time in this study.

The prospects of Fed-batch fermentation in brewing

The study has only evaluated the performance of fed-batch to batch fermentation without supplementation looking only at temperature, as it is significant in

fermentation. More research should look at the effect of other factors, such as yeast concentration, pH, *etc.*

Spoilage test of wort material should also be conducted during production to check the practicality of implementing fed-batch at high volume commercial productions. The test should include evaluation to investigate if stored wort that is not used instantly in fed-batch will have any negative effect on the taste of the final beer. In addition, this study recommends further work to quantify additional costs that might be incurred industrially to implement fed-batch fermentation process. Detailed economic benefits feasibility gains from using fed-batch fermentation will encourage industry to embrace change in their operation.

Heat integration for batch processes

In our study of optimisation of brewing (batch) process using integration, energy minimization was used as the only objective function. In industry, there might be several production objectives such as multi-product needs, production quotas and deadlines, complex production plans, and resource limitations. For these reasons, in other industries, energy recovery might be secondary objective. It is thus recommended that economic, technical and business decisions be made to evaluate if energy minimization outweighs other objectives when setting up process integration.

REFERENCES

Abbasi, B., Mahlooji, H. (2012). Improving response surface methodology by using artificial neural network and simulated annealing, *Expert Systems with Applications*, vol. 39, no. 3, 3461-3468.

Adams, W., James, B. (1995). *The Structure of American Industry*, 9th ed., Englewood Cliffs, New Jersey: Prentice Hall.

Adeyemo, J., Eniton, A. (2011). Optimisation of Fermentation Process using Evolutionary Algorithms, *Journal of Bioengineering*, vol. 16, 1464-1472.

Adonyi, R., Romero, J., Puigjaner, L., Friedler, F. (2003). Incorporating Heat Integration in Batch Process Scheduling, *Applied Thermal Engineering*, vol. 23, 1743-1762.

Ahmad, S., Hui, D. C. W. (1991). Heat recovery between areas of integrity, *Computers and Chemical Engineering*, vol. 15, no. 12, 809-832.

Anderson, R. G. (2003). In *Alcohol and Temperance in Modern History: An International Encyclopaedia*, Blocker, J., Tyrrell, I., Fahey, D. (Editors). vol.1, ABC-Clio, Santa Barbara.

Ashton, G. J, Obeng, E. D. A. (1987). Process integration of batch processes, In: 79th AIChE Annual Meeting, New York.

Aslan, N. (2007). Application of response surface methodology and central composite rotatable design for modelling the influence of some operating

variables of a Multi-Gravity Separator for coal cleaning, *Fuel*, vol. 86, no.5–6, 769-776.

Baishen, F., Hongwen, C., Ning, W., Zongding, H. (2003). Using Genetic Algorithms coupling Neural Networks in a study of Xylitol production: Medium Optimisation, *Process Biochemistry Journal*, vol. 38, 979-985.

Baker, J., Grefenstette, J. (1989). How genetic algorithms work: a critical look at implicit parallelism. *Proceedings of the Third International Conference on Genetic Algorithms*, 20-27.

Bamforth, C. (2003). Beer: Tap into the art and science of brewing, 1st and 2nd ed., Oxford University Press, Inc.

Bas, D., Boyaci, H. (2007). Modeling and optimisation I: Usability of response surface methodology, *Journal of food engineering*, 78, 836-845.

Batunde, P. (2007). Optimisation of the Fermentation in a Brewery with Fuzzy logic controller, *Biotechnology Series*, vol. 11, 79-92.

Baxter, E. D., Hughes P. S. (2001). Beer: quality, safety and nutritional aspects, Royal Society of Chemistry, Cambridge.

Beale, R., Jackson, T. (1990). Neural Computing: an introduction, Bristol: Adam Hilger, Philadelphia and New York.

Becker, T., Enders, T. Delgado, A. (2002). Dynamic Neural Network as a tool for online Optimisation of Industrial Fermentation, *Journal of Bioprocessing Engineering*, vol. 24, 347-354.

- Bhatia, T., Biegler, L. T. (1996). Dynamic optimisation in the design and scheduling of multiproduct batch plants, *Industrial and Engineering Chemistry Research*, vol. 35, 2234-2246.
- Bhatia S, Wong C. T, Abdullah A. Z. (2009). Optimisation of air-borne butyl acetate adsorption on dual-function Ag-Y adsorbent-catalyst using response surface methodology, *Journal of Hazard Materials*, vol. 164, 1110-1117.
- Bhotmange, M., Shastri, P. (2011). Application of Artificial Neural Networks to Food and Fermentation Technology, *Artificial Neural Networks - Industrial and Control Engineering Applications*, 1-10.
- Bishop, C. M. (1995). *Neural Networks for Pattern Recognition*, Oxford: Claredon Press, New York.
- Boivin, B., Malanda, M. (1997). Improvement of malt quality and safety by adding starter culture during the malting process, *Technical Quarterly Master Brewers Association of Americas*, vo. 34, 358–363.
- Boswell, C. D., Nienow, A. W., Gill, N. K., Kocharunchitt, S., Hewitt, C. J. (2003). The Impact of Fluid Mechanical Stress on *Saccharomyces Cerevisiae* Cells During Continuous Cultivation in an Agitated, Aerated Bioreactor; its implication for Mixing in the Brewing Process and Aerobic Fermentations, *Food and Bioproducts Processing*, vol. 81, no. 1, 23-32.
- Box, G. E. P., Hunter, J. S. (1957). Multi-factor experimental design for exploring response surfaces, *Annals of Mathematical Statistics*, vol.28, 195-241.

Bozan, M., Borak, F. I. (2001). A computerized and integrated approach for heat exchanger network design in multipurpose batch plants, *Chemical Engineering Processes*, vol. 40, 511–24.

Bramlette, M. F., Bouchard, E. E. (1991). Genetic algorithms in parametric design of aircraft: Handbook of Genetic Algorithms, Van Nostrand Reinhold, New York.

Briggs, D. E., Hough, J. S., Stevens, R., Young, T. W. (1981a). Malting and Brewing Science, 2nd ed., vol. 1, Chapman and Hall, New York.

Briggs, D.E. (1998). Malts and Malting, Kluwer Academic / Plenum Publishers, New York.

Cheema, J. J. S., Sankpal, N. V., Tambe, S. S., Kulkarni, B. D. (2002). Genetic programming assisted stochastic optimisation strategies for optimisation of glucose to gluconic acid fermentation. *Biotechnology Progress*, vol. 18, 1356–1365.

Choi, D. B., Park, E. Y. (2006). Enhanced production of mouse alpha amylase by feeding combined nitrogen and carbon sources in fed-batch culture of recombinant *Pichia pastoris*, *Process & Biochemical Engineering*, vol.41, no.2, 390–397.

Cinar, A., Parulekar, S. J., Undey, C., Birol, G. (2003). Batch Fermentation: Modeling, Monitoring, and Control. Marcel Dekker, Inc. New York.

Colombie, V., Bideaux, C., Goma G., Uribe Larrea, J. L. (2005). Effects of glucose limitation on biomass and spiramycin production by *Streptomyces ambofaciens*. *Bioprocess Engineering*, vol. 28, no.1, 55–61.

Comrie, A. C. (1997). Comparing Neural Networks and Regression models for ozone forecasting, *Journal of Air and Waste Management Association*, vol. 47, 653-663.

Cook, H., Skadron, K. (2007). Genetically Programmed Response analysis for efficient design space exploration: Technical Report, University of Virginia, USA.

Corominas, J., Espuña, A., Puigjaner, L. (1994). Method to incorporate energy integration considerations in multiproduct batch processes, *Computers in Chemical Engineering*, vol.18, 1043-1055.

Cowpe, J. S., Astin, J. S., Pilkington, R. D. Hill, A. E. (2007). Application of Response Surface Methodology to laser-induced breakdown spectroscopy: Influences of hardware configuration, *Spectrochimica Acta Part B: Atomic Spectroscopy*, vol. 62, no. 12, 1335-1342.

Cybenko, G. (1989). Approximation by Superposition of a Sigmoidal Function, *Mathematical Control Signal Systems*, vol. 2, 303-314.

Dasari, V. R., Donthireddy, R. S., Nikko, Y.M., Garapti, H. R. (2009). Optimisation of medium constituents for Cephalosporin production using response surface methodology and Artificial Neural Networks, *Journal of Biochemical Technology*, vol. 3, 69-74.

De Coninck, J., Bouquelet, S., Dumortier, V., Duyme, F., Denantes, V. I. (2000) Industrial media and fermentation process for improved growth and protease production by *Tetrahymena thermophila* BIII, *Journal of Industrial Microbiology and Biotechnology*, vol. 24, 285–290

Demirci, A., Pometto III, A. L., Cox, D. J. (1999). Enhanced organically bound selenium yeast production by fed-batch fermentation, *The Journal for Agricultural Food*, vol. 47, no. 6, 2496–2500.

Dombek K, M., Ingram L, O. (1987). Ethanol Production during Batch Fermentation with *Saccharomyces Cerevisiae*: Changes in Glycolytic Enzymes and Internal pH, Department of Microbiology and Cell Science, University of Florida, Florida.

Dragone, G., Mussatto, S. I., Almeida e Silva, J. B. (2008). Influence of temperature on continuous high gravity brewing with yeasts immobilized on spent grains, *The Journal of European Food and Research Technology*, vol. 228, 257-264.

Draper, N. R., Lin, D. K. (1990). Small Response-Surface Designs, *Technometrics*, vol. 32, no. 2, 187-194.

Essan, N., Zita, B., Emmanuel, A. N., Patrice, K., Ismael, D., Benjamin, Y. (2009). Modelling of Osmotic Dehydration of Mango (*Mangifera Indica*) by Recurrent Artificial Neural Network and Experimental Design, *Research Journal of Agriculture and Biological Sciences*, vol. 5, no. 5, 754-761, 2009

Farzad, D., Faezeh, N. (2011). Modeling and optimisation of the bi-substrate peroxidase-enzyme catalyzed potentiometric assay of hydrogen peroxide by response surface methodology with a central composite rotatable design, *Journal of Molecular Catalysis B: Enzymatic*, vol. 68, no. 2, 162-167.

- Fernandez, I., Renedo, C. J., Perez, S. F., Ortiz, A., Manana, M. (2012). Energy recovery in batch processes, *Renewable and Sustainable Energy Reviews*, vol. 16, 2260– 227.
- Floudas, C. A., Lin, X. (2005). Mixed Integer Linear Programming in Process Scheduling: Modeling, Algorithms, and Applications, *Annals of Operations Research*, vol. 139, 131–162.
- Gardner, M. W., Dorling, S. R. (1998). Artificial Neural Networks (The Multilayer Perceptron), *Journal of Atmospheric Environment*, vol. 32, 2627-2636.
- Goldberg, D. E. (1989). Genetic algorithms in searching, optimisations and machine learning, Addison-Wesley, Boston MA.
- Goldhammer, T. (2008). The Brewer's Handbook, 2nd ed., Apex, USA.
- Gopal, C. V., Hammond, J. R. M., Pye, J. M. (1993). Interactive control of beer fermentation: Demonstration at the laboratory scale, *Masters Brewers Association of Americas Technology Quarterly*, vol. 30, 71-79.
- Gors ek, A., Glavic, P. (1997). Design of batch versus continuous processes Part I: Single-Purpose Equipment, *Chemical Engineering Research and Design*, vol. 75, no. 7, 709-717.
- Gregersen, L., Jorgensen, S. B. (1999). Supervision of fed-batch fermentation, *Chemical Engineering Journal*, vol. 75, 69-76.
- Guyton, C. A., Hall, J. E. (2006). Textbook of Medical Physiology, Elsevier Inc, Philadelphia.

- Haaland, P. D. (1989). Statistical problem solving, In: Haaland, P. D (Editor), Experimental design in biotechnology. Marcel Dekker Inc., New York, 1–18.
- Halim, I., Srinivasan, R. (2009). Sequential Methodology for Scheduling of Heat-Integrated Batch Plants, *Industrial and Engineering Chemistry Research*, vol. 48, 8551–8565.
- Hanrahan, G. (2011). Artificial Neural Networks in Biological and Environmental Analysis, CRC Press, New York.
- Harrison, M. A. (2009). Beer/Brewing, Encyclopaedia of Microbiology, 3rd ed., Elsevier Inc, Athens, GA, USA, 23-33.
- Hewitt, C. J., Nienow, A. W. (2007). The scale-up of microbial batch and fed-batch fermentation processes, *Advances in Applied Microbiology*, no. 62, 105-135.
- Holland, J. H. (1992). Genetic Algorithms, *Scientific American*, 44-50.
- Hornik, K., Stichcombe, M., White, H. (1989). Multilayer layer feed forward networks are universal approximators, *Neural Networks*, vol. 2, 359-366.
- Hornsey, I. S. (1999). Brewing, Nethergate Brewery Co Ltd, New York.
- Hornsey, I. S. (2004). A History of Beer and Brewing, 1st ed., Royal Society of Chemistry, Washington D.C.
- Houck, C., Joines, J. (1995). A Genetic Algorithm for Function Optimisation: A Matlab Implementation, Technical Report: North Carolina State University, USA.
- Hough, J. S. (1985). The Biotechnology of malting and brewing, Cambridge University Press, New York.

- Huang, W., Chen, B. (2007). Scheduling of batch plants: constraint-based approach and performance investigation, *International Journal of Production Economics*, vol. 105, 425–44.
- Huang, Y., Kangas, L. J., Rasco, A. B. (2007). Application of Artificial Neural Networks (ANN) in Food Science, Tay and Francis Publishers, New York.
- Hunt, K. J., Sbarbaro, D., Zbikowski, R., Gawthrop, P. J. (1992). Neural Networks for control systems – A survey, *Automatica*, vol. 28, 1083-1099.
- Ierapetritou, M. G., Floudas, C. A. (1998). Effective Continuous-Time Formulation for Short-Term Scheduling 1: Multipurpose Batch Processes, *Industrial Engineering and Chemistry Research*, vol. 37, no. 11, 4341-4359.
- Iyuke, S. E., Madigoe, E. M., Maponya, R. (2008). The Effect of Hydroxycinnamic Acids and Volatile Phenols on Beer Quality, *Journal of the Institute of Brewing*, vol. 114, no. 4, 300-305.
- Jain, S. K., Singh, V. P. (2003). Water Resources Systems Planning and Management, Elsevier, Amsterdam, 856-858.
- John, P. A. (2005). Origin and History of Beer and Brewing: From Prehistoric Times to the Beginning of Brewing Science and Technology, BeerBooks, Cleveland, Ohio, pp 34.
- Jong, J., Hsiun, D., Wu, W. (1995). Fed-Batch Culture of *Bacillus thuringiensis* for Thuringiensin Production in a Tower Type Bioreactor, *Biotechnology and Bioengineering*, vol. 48, 207-213.

Johnson, C.R., Burnham, K.J, James, D.J.G. (1998). Extended simulation model for a brewery fermentation process. *International Conference on Simulation*, pp: 259-262.

Jung, S., Lee, T .I, Yang, D. R., Chang K. S. (1994). Synthesis of maximum energy recovery networks in batch processes, *Korean Journal of Chemical Engineering*, vol. 11, no. 3, 162-171.

Keller, G. (2001). *Applied statistics with Microsoft Excel*, Duxbury, California.

Kemp, I. C., Macdonald, E. K. (1988). Application of Pinch Technology to Separation, Reaction and Batch Processes, *Industrial Chemical Engineering Symposium, Ser. No. 109, Understanding Process Integration II*, 239-258.

Kemp, I. C., Deakin, A. W. (1989a). The Cascade Analysis for Energy and Process Integration of Batch Processes Part 1: Calculation of Energy Targets, *Chemical Engineering Research and Design*. vol. 67, 495-509.

Kemp, I. C., Deakin, A. W. (1989b). The cascade analysis for energy and process integration of batch processes. Part 2: network design and process scheduling, *Chemical Engineering Research and Design*, vol. 67, 510-516.

Kemp, I. C., Deakin, A. W. (1989c). The cascade analysis for energy and process integration of batch processes. Part 3: a case study, *Chemical Engineering Research and Design*, no. 67, 517-525.

Khoo, L. P., Chen, C. H. (2000). Integration of Response Surface Methodology with Genetic Algorithms, *Journal of Advanced Manufacturing Technology*, vol. 18, 148-489.

Khuri, A. I., Cornell, J. A. (1987). Response Surfaces: Designs and Analyses. Marcel Dekker Inc, New York.

Kiran, M. D., Shrikant, A. S., Parag, S. S., Lele, S. S., and Rekha, S.S. (2008). Comparison of artificial neural network (ANN) and response surface methodology (RSM) in fermentation media optimisation: Case study of fermentative production of scleroglucan, *Biochemical Engineering Journal*, vol. 41, no. 3, 266-273.

Koc, A. B., Heinemann, P. H., Ziegler, Z. G. (2007). Optimisation of Whole Milk Powder processing variables with Neural Networks and Genetic Algorithms, *Food and Bioproducts Processing : Transactions of the Institution of Chemical Engineers, Part C*, vol. 85, 336-343.

Koehn, P. (1994). Combining Genetic Algorithms and Neural Networks. University of Tennessee Press, Knoxville

Kondili, E., Pantelides, C. C., Sargent, R. W. H. (1993). A general algorithm for short-term scheduling of batch operations: I, MILP formulation. *Computers & Chemical Engineering*, vol. 17, no. 2, 211-227.

Kotecha, P. R., Kapadi, M. D. (2003). Multi-Objective Optimisation Issues in Short-Term Batch Scheduling, *Proceedings of the 17th World Congress, The International Federation of Automatic Control*, Seoul, Korea.

Krishnakumar, K., Goldberg, D. E. (1992). Control system optimisation using genetic algorithms, *Journal of Guidance, Control, and Dynamics*, vol. 15, no.3, 735-740.

Kumar, M. S., Jana, S. K., Senthil, V., Shashanka, V., Vijay, S., Sadhukhan, K. A. K. (2000). Repeated fed-batch process for improving lovastatin production, *Process Biochemistry*, vol. 36, no. 4, 363-368.

Lee, B., Reklaitis, G. V. (1995a). Optimal Scheduling of a Cyclic Batch Processes for Heat Integration - I. Basic Formulation, *Computers and Chemical Engineering*, vol. 19, 883-905.

Lee, B., Reklaitis, G. V. (1995b). Optimal scheduling of cyclic batch processes for heat integration. Part II. Extended problems, *Computers and Chemical Engineering*, vol.19, 907-31.

Lee, J., Lee, S. Y., Park, S., Middelberg, A. P. J. (1999). Control of fed-batch fermentations, *Biotechnology Advances*, vol. 17, no. 1, 29 - 48.

Linko, S. (1998a). Expert systems- what they can do for food industry?, *Trends Food Science Technology*, vol. 9, 3-12.

Linko, M., Haikara, A., Ritala, A., Penttilä, M. (1998). Recent advances in the malting and brewing industry , *Journal of Biotechnology*, vol. 65, no. 2-3, 85-98.

Linnhoff, B., Townsend, D.W., Boland, D., Hewitt, G.F., B.E.A. Thomas, A.R. Guy, Marsland, R. H. (1982). A User Guide to Process Integration for the Efficient Use of Energy, last edition, IChemE, Rugby, UK.

Liu, C. H., Hwang, C. F., Liao, C. C. (1999a). Medium optimisation for glutathione production by *Saccharomyces Cerevisiae*, *Process Biochemistry*, vol. 34, 17–23.

Liu, T. J., Miura, S., Arimura, T., Tei, M. Y., Park, E. Y., Okabe, M. (2005). Evaluation of L-lactic acid production in batch, fed-batch, and continuous cultures of *Rhizopus* sp MK-96 1196 using an airlift bioreactor, *Biotechnology Bioprocess Engineering*, vol. 10, no. 6, 522–527.

Longobardi, G. P. (1994). Fed-batch vs. batch fermentation, *Bioprocess Engineering*, vol. 10, 1985-194.

Madigoe, E. M. (2009). Commissioning and Optimisation of Wits Micro-brewery plant, MSc dissertation, University of the Witwatersrand, Johannesburg, South Africa.

Majozi, T. (2006). Heat Integration of Multipurpose Batch Plants Using a Continuous-time Framework, *Applied Thermal Engineering*, vol. 26, 1369-1377.

Manzano, M., Giusto, C., Bartolomeoli, I., Buiatti, S., Comi, G. (2005). Microbiological analyses of dry and slurry yeasts for brewing, *Journal Institute of Brewing*, vol. 111, no. 2, 203-208.

Marchitan, N., Cojocaru, C., Mereuta, A., Duca, G., Cretescu, I., Gonta, M. Modeling and optimisation of tartaric acid reactive extraction from aqueous solutions: A comparison between response surface methodology and artificial neural network, *Separation and Purification Technology*, vol. 75, no. 3, 273-285.

McInerney, M., Dhavan, A. P. (1993). Use of Genetic Algorithms with Backpropagation in Training of Feed-Forward Neural Networks, *International Conference of Neural Networks*, vol. I, 203-208.

Mendez, C. A., Cerda, J., Grossmann, I. E., Harjunkoski, I., Fahl, M. (2006). State-of-the-art review of optimisation methods for short-term scheduling of batch processes, *Computers and Chemical Engineering*, vol. 30, no.6-7, 913-946.

Michael, N. A. E., Shahidi, F. (1990). Biochemistry of Foods, 2nd ed., Elsevier Inc, USA, 297-334.

Michalewicz, Z. (1996). Genetic Algorithms + Data Structures = Evolution Programs, Springer-Verlag, Berlin Heidelberg,

Muster-Slawitsch, B, Weiss, W., Schnitzer, H., Brunner, C. (2011). The green brewery concept- energy efficiency and the use of renewable energy sources in breweries, *Journal of Applied Thermal Engineering*, vol. 31, no. 13, 2123-2134.

Myers, R. H., Montgomery, D. C. (2002). Response Surface Methodology: Process and product optimisation using designed experiments, 2nd ed., John Wiley and Sons Inc, New York.

Nagata, K., Khim, H. (2003). Optimisation of fermentation medium using neural networks and genetic algorithms, *Journal of Applied Biotechnology*, vol. 25, 1837-1842.

Neubauer, P., Ruottinen, M. (2008). Fed-Batch Technology: Theory of Fed-batch. *BioSilta*, 4 (8).

Nilsson, A., Taherzadeh, M. J., Lidén, G. (2001). Use of dynamic step response for control of fed-batch conversion of lignocellulose hydrolyzates to

ethanol, *Journal of Biotechnology*, vol. 89, 41–53.

Niu, D., Wang, F., Zhang, L., He, D., Jia, M. (2011). Neural network ensemble modeling for nosiheptide fermentation process based on partial least squares regression, *Chemometrics and Intelligent Laboratory Systems*, vol. 105, no. 1, 125-130.

Obeng E. D. A., Ashton G. J. (1988). On pinch technology based procedures for the design of batch process, *Chemical Engineering Research and Design*, vol. 66, 255–259.

Obeng, D. P., Morrel, S., Napier-Munn, T. J. (2005). Application of central composite design to modeling the effect of some operating variables on the performance of the three product cyclone, *International Journal of Mineral Processing*, vol. 76, 181-192.

Oraon, B., Majumdar, G., Ghosh, B. (2006). Application of response surface method for predicting electroless nickel plating, *Materials and Design*, vol. 27, 1035–1045.

Papageorgiou, L. G., Shah, N., Pantelides, C. C. (1994). Optimal Scheduling of Heat-Integrated Multipurpose Plants, *Industrial Engineering and Chemistry Research*, vol. 33, 3168.

Parsons, R., Johnson, M. E. (1997). Case Study In Experimental Design Applied To Genetic Algorithms With Applications To DNA Sequence Assembly, *American Journal of Mathematical and Management Sciences*, vol. 17, 369-396.

Patan, K. (2008). Artificial Neural Networks for the Modelling and Fault Diagnosis of Technical Processes, Springer- Verlag, Berlin Heidelberg.

Phisalaphong, M., Srirattana, N., Tanthapanichakoon, W. (2006). Mathematical modelling to investigate temperature effect on kinetic parameters of ethanol fermentation, *Biochemical Engineering Journal*, vol. 28, no. 1, 36-43.

Pinto, T., Novais, A. Q., Barbosa-Povoa, A. P. F. D. (2003). Optimal Design of Heat-integrated Multipurpose Batch Facilities with Economic Savings in Utilities: A Mixed Integer Mathematical Formulation, *Annals of Operations Research*, vol. 120, 201-230.

Pollanen, J., Rousu, J., (2001) *A neural network tool for brewery fermentation*, <http://www.cs.helsinki.fi/>, Accessed May 2012.

Pourali, O., Amidpour, M., Rashtc, D. (2006). Time decomposition in batch process integration, *Chemical Engineering and Process Intensification*, vol. 45, no. 1, 14–21.

Powell, C. D., Quain, D. E., Smart, K. A. (2003). The impact of brewing yeast cell age on fermentation performance, attenuation and flocculation, *FEMS Yeast Research*, vol. 3, no. 2, 149-157.

Pramanik, K. (2004). Use of artificial neural networks for prediction of cell mass and ethanol concentration in batch fermentation using *Saccharomyces Cerevisiae* yeast, *IE (I) Journal-CH*, vol. 85, 31-35.

Priest, F. G., Stewart, G. (2006). Handbook of Brewing, 2nd ed., CRC press, London, New York.

- Puigjaner, L. (2007). Extended modeling framework for heat and power integration in batch and semi-continuous processes, *Chemical Product and Process Modeling*, vol. 2, 1–46.
- Rajkumar, T., Bardina, J. (2005). Training data requirement for a neural network to predict aerodynamic coefficients, NASA Research Center, USA.
- Ranganath, M., Rengannathan, S., Gokulnath, C. (2002). Identification of bioprocesses using genetic algorithms, *Bioprocesses Engineering*, vol. 24, 123-127.
- Reed, G., Nagodawithana, T. W. (1991). Yeast Technology, 2nd ed., Van Nostrand Reinhold, New York.
- Rumelhart, D. E. (1986). Learning internal representations by Error Propagation, MIT Press, Massachusetts.
- Runquist, D., Hahn-Hägerdal, B., Bettiga, M. (2009). Increased expression of the oxidative pentose phosphate pathway and gluconeogenesis in anaerobically growing xyloseutilizing *Saccharomyces Cerevisiae*. *Microbial Cell Factories*, vol. 49, no. 8, [doi: 10.1186/1475-2859-8-49](https://doi.org/10.1186/1475-2859-8-49).
- Saarela, U., Leiviskä, K., Juuso, E. (2003). Modelling of a Fed-Batch Fermentation Process, Control Engineering Laboratory, Department of Process and Environmental Engineering, University of Oulu, Finland.
- Sadr-Kazemi, N., Polley G. T. (1996). Design of energy storage system for batch process plants, *Trans. IChemE Part A*, vol. 74, 584–596.

Shimizu, H., Miura K., Shioya, S., Suga, K. (1993). An overview on the control system design of bioreactors, *Advances in Biochemical Engineering and Biotechnology*, vol. 50, 65-84.

Shimizu, K. (1997). Bioprocess systems analysis method-the principle of systems analyses method and its application. Corona Publishing Co. Ltd, Tokyo, Japan, pp. 225.

Shuler, M., Kargi, F. (2005). Bioprocess Engineering: Basic Concepts 2nd ed., Case, Prentice-Hal, New Jersey.

Simate, G. S. (2009). Bacterial Leaching Of Nickel Literites Using Chemolithotrophic Micro-Organism, MSc Dissertation, University of the Witwatersrand, Johannesburg, South Africa.

Stolze, S., Mikkelsen, J., Lorentzen, B., Petersen, P. M., Qvale, B. (1995). Waste-heat recovery in batch processes using heat storage, *Journal of Energy Resource Technology*, vol. 117, 142–149.

Syu, M. J., Tsao, G.T. (2010). Neural Networks modelling of batch cell growth pattern, *Biotechnology Engineering*, vol. 42, 376-380.

Tokos, H., Pintarič, Z. N., Glavič, P. (2009). Energy Saving Opportunities in Heat Integrated Beverage Plant Retrofit, *Applied Thermal Engineering*, vol. 30, no. 1, 36-44.

Trelea, I. C., Titica, M., Corrieu, G. (2004). Dynamic optimisation of the aroma production in brewing fermentation, *Journal of Process Control*, vol. 14, 1-16.

- Triadaphillou, S., Martin, E., Montague, G., Jeffkins, P., Stimpson, S., Nordon, A. (2005). Enhanced modeling of an industrial fermentation process through data fusion techniques , *Computer Aided Chemical Engineering*, vol. 20, 1393-1398.
- Tsai, J. S. (1993). Integration of Genetic Algorithms to Engineering Optimisation Problems, Master's Thesis, Ohio State University, Ohio.
- Vaselenak, J. A., Grossmann, I. E., Westerberg, A. W. (1986). Heat integration in batch processing, *Industrial and Engineering Chemistry Process Design and Development*, vol. 25, 357- 366.
- Vassileva, S., Mileva, S. (2010). AI-based software tools for beer brewing monitoring and control, *Biotechnology and Biotechnological Equipment*, vol. 24, doi: 10.2478/v10133-010-0060-0.
- Venkata, R. K., Reddy, D., Murali, Y. N., Hanumantha, G. (2009). Optimisation of medium constituents for Cephalosporin C production using Response Surface Methodology and Artificial Neural Networks, *Journal of Biochemical Technology*, vol. 1, no. 3, 69-74.
- Vlassides, S., Ferrier, G. J., Block, D. E. (2000). Using Historical Data for Bioprocess Optimisation: Modelling Wine Characteristics Using Artificial Neural Networks and Archived Process Information, *Biotechnology and Bioengineering Journal*, vol. 73, 55-68.
- Vu, T. K. L., Le, V. V. M. (2007a). Application of fed-batch fermentation in high-gravity brewing, *Journal of Science - Natural Sciences and Technology*, vol. 23, 166– 173.

- Vu, T. K. L., Le, V. V. M. (2010). Using fed-batch fermentation in high-gravity brewing: effects of nutritional supplementation on yeast fermentation performance, *International Food Research Journal*, vol.17, 117-126.
- Wang, Y. P., Smith, R. (1995). Wastewater Minimization with low rate constraints, *Trans. IChemE*, vol. 73, Part A, 889-904.
- Whitley, D. (1995). Genetic Algorithms in Engineering and Computer Science., John Wiley & Sons Ltd., London.
- Wilson, R. G., Gourvish, T. (1993). The production and consumption of alcoholic beverages in Western Europe, *Alcoholic Beverages and European Society*. Annex 1. The Historical, Cultural and Social Roles of Alcoholic Beverages, The Amsterdam Group, Eds., Elsevier, Amsterdam.
- Windrow, B., Lehr, M. A. (1990). 30 Years of adaptive neural networks: Perception, madaline and backpropagation, *Proceedings of the IEEE*, vol. 78, 1441-1457.
- Wu, S., Yu, X., Hu, Z., Zhang, L., Chen, J. (2009). Optimizing aerobic biodegradation of dichloromethane using response surface methodology, *Journal of Environmental Science (China)*, vol. 21, 1279-1283.
- Ya Ling, Z., Morimura, S., Kida, K. (1995). Effect of fermentation temperature on relationship between cell viability and trehalose content of *Saccharomyces cerevisiae* KF-7 in repeated-batch fermentation, *Journal of Fermentation and Bioengineering*, vol. 80, no. 2, 204-207.

- Ye, K., Jin, S., Shimizu, K., (1996). Performance improvement of lactic acid fermentation by multistage extractive fermentation, *Journal of fermentation and Bioengineering*, vol. 81, no. 3, 240-246.
- Yi-Huang, C., Ku-Shang, C., Cheng-Wei, H., Chuan-Liang, H., Hung-Der, J. (2012). Comparison of batch and fed-batch fermentations using corncob hydrolysate for bioethanol production, *Fuel*, vol. 97, 166-17.
- Youssef, C. B., Guillou, V., Olmos-Dichara, A. (2000). Modelling and adaptive control strategy in a lactic fermentation process, *Control Engineering Practice*, vol. 8, no. 11, 1297-1307.
- Zealand, C. M., Burn, D. H., Siminovic, S. P. (1999). Short term stream flow forecasting using artificial neural networks, *Journal of Hydrology*, vol. 214, 32-48.
- Zhang, X., Liu, J., Qiao, H., Liu, H., Ni, J., Zhang, W., Shi, Y. (2010). Formulation optimisation of dihydroartemisinin nanostructured lipid carrier using response surface methodology, *Powder Technology*, vol. 197, no. 1–2, 120-128.
- Zhao, X. G., O'Neill, B. K., Roach, J. R. (1998). Wood RM. Heat integration for batch processes. Part 1. Process scheduling based on cascade analysis, *Chemical Engineering Research Design*, vol. 76, 685–99.
- Zhong, Y. L., Morimura, S., Kida, K. (1995). Effect of fermentation temperature on relationship between cell viability and trehalose content of *Saccharomyces Cerevisiae* KF-7 in repeated-batch fermentation, *Journal of Fermentation and Bioengineering*, vol. 80, no.2, 204-207.

Zhu, Y., Zhang, J., Shi, Z., Mao, Z. (2004). Optimisation of Operating conditions in the Rice heat blast process for Chinese Rice Wine Production by Computational Utilisation of Neural Network and Genetic Algorithms, *Journal of the Institute of Brewing and Distilling*, vol. 110, no. 2, 117-123.

Zou, T., Cai, M., Du, R., Liu, J. (2012). Analyzing the uncertainty of simulation results in accident reconstruction with Response Surface Methodology, *Forensic Science International*, vol. 216, no.1–3, 49-60.

APPENDIX A: RSM MODEL VALIDATION

Total Sum of Squares

The total variation in a set of data is called the Total Sum of Squares (SS_T). It is computed by summing the squares of the deviations of the observed $y_{i,exp}$'s about

their average value, $\bar{y} = \frac{(y_1 + y_2 + \dots + y_n)}{N}$ and then using the equation given below (Khuri and Cornell, 1987).

$$SS_T = \sum_{i=1}^N (y_{i,exp} - \bar{y})^2 \quad (A1)$$

The quantity SS_T is associated with $N - 1$ degrees of freedom. The square term is used as a means of factoring out whether the sample is below or above the mean. If the square is not taken, variations below the mean may cancel variation above the mean to make it seem as if there is no variation (Keller, 2001). The total sum of squares can be partitioned into two parts; the sum of squares due to regression (sum of squares explained by the fitted model) and the sum of squares unaccounted for by the fitted model.

Sum of Squares of Regression

The formula for calculating the sum of squares due to regression (SS_R) is (Khuri and Cornell, 1987):

$$SS_R = \sum_{i=1}^N (y_{i,cal} - \bar{y})^2 \quad (A2)$$

The deviation $y_{i,cal}$ is the difference between the value predicted by the model for the i^{th} observation and the overall average of the observed, $y_{i,exp}$'s. If the model contains p -parameters, then the degree of freedom associated with SS_R is $p-1$.

Sum of Squares of Errors

The sum of squares of errors (SSE) unaccounted for by the fitted model or sum of squares of residuals or error is computed as follows (Khuri & Cornell, 1987):

$$SS_R = \sum_{i=1}^N (y_{i,exp} - y_{i,cal})^2 \quad (A3)$$

The deviation, $y_{i,exp}-y_{i,cal}$ is the difference between the observed value and value predicted by the model for the i th observation. The number of degrees of freedom of SS_E is $N-p$.

Determination of SS_T , SS_R , and SS_E using matrices

Using matrix notation and letting O be a $1 \times N$ vector of ones, the quickest methods for calculating SS_T , SS_R , and SS_E are given as (Khuri and Cornell, 1987):

$$\begin{aligned} SS_T &= y^T y - \frac{(Oy)^2}{N} \\ SS_R &= \beta^T X^T y - \frac{(Oy)^2}{N} \\ SS_E &= y^T y - \beta^T X^T y \end{aligned} \quad (A4)$$

Where y is an $N \times 1$ vector of the observations; X is an $N \times (k+1)$ matrix of the levels of the independent variables; β is a $(k+1) \times 1$ vector of the estimates of the

regression coefficients; y^T , X^T and β^T refers to the transpose of y , X and β , respectively.

Test of lack of fit of the fitted model

A procedure for checking the adequacy of the fitted model is called testing lack of fit of the fitted model (Khuri and Cornell, 1987). The fitted model is inadequate or is lacking in fit when it does not contain sufficient number of terms due to the omission of factors (other than those in the proposed model) that affect the response, and/or, omission of higher-order terms involving the factors in the proposed model. The inability of the fitted model to adequately account for the variation in the observed response values is reflected in the portion of the total variation that is called residual variation or the variation unaccounted for by the fitted model (Khuri and Cornell, 1987, Myers and Montgomery, 2002).

Isolating the portion of the residual variation that is directly attributed to under fitting the true surface with the fitted model is necessary in order to test for adequacy of the fitted model. The residual sum of squares, SS_E , can be partitioned into two sources of variation; (1) the variation among the replicates at the design points where replicates have been collected, and (2) variations arising from lack of fit of the fitted model.

The sum of squares due to replicate observations is called the sum of squares due to pure error (SS_{PE}). In other words, it is the sum of squares of differences between all the individual experimental values and the average of the experimental values at the same level. Once SS_{PE} is calculated, it is then subtracted from the SS_E to produce the sum of squares due to lack of fit (SS_{LOF}) (Simate, 2009). SS_{PE} can also be calculated from the replicated centre points of the experiments (Myers and Montgomery, 2002) as given below:

$$SS_{PE} = \sum_{\text{centre-points}} (Y_{c,i} - \bar{Y}_c)^2 \quad (\text{A5})$$

where $Y_{c,i}$ are the observations of the centre points, $\bar{Y}_c$ is the average of the observations of the n_c runs at the centre. The number of degrees of freedom associated with SS_{PE} in this case is $n_c - 1$. By subtraction, the sum of squares due to lack of fit is then calculated as (Khuri and Cornell, 1987):

$$SS_{LOF} = SS_E - SS_{PE} \quad (\text{A6})$$

The test of the null hypothesis of adequacy of fit (or lack of fit is zero) involves calculating the value of the Fisher's variance ratio test (F -test):

$$F = \frac{SS_{LOF} / (N - p - n_c + 1)}{SS_{PE} / (n_c - 1)} \quad (\text{A7})$$

where n_c is the number of replicates of centre points, p is the number of terms in the fitted model; N is the number of observations. $N - p - n_c + 1$ is the degree of freedom for SS_{LOF} .

The calculated value is then compared with the standard table value of F ratio for the desired level of confidence; α . Lack of fit exists at α level of significance if the calculated value of F exceeds the standard table value, $F_{\alpha, N-p-n_c+1, n_c-1}$ (Khuri and Cornell, 1987). If not, the model can be accepted at the desired confidence level as providing an adequate representation of the data.

Test for significance of regression

The test for significance of regression is a test to determine whether a linear relationship exists between the response variable, y , and the regression coefficients, β_i (Myers and Montgomery, 2002). The usual test of the significance of the fitted regression equation is a test of the null hypothesis, H_0 : all values of β_i (excluding β_0) are zero. The alternative hypothesis is, H_a : at least one value of β_i (excluding β_0) is not zero. Assuming normality of the errors, the test of H_0 involves first calculating the value of the F-statistic,

$$F = \frac{SS_R / (p-1)}{SS_E / (N-p)} \quad (A8)$$

If the null hypothesis is true, then the F -statistic Equation (A8) follows an F-distribution with $p-1$ and $N-p$ degrees of freedom under the ν_1 (numerator) and ν_2 (denominator) F -distribution, respectively, where p is the number of terms in the fitted model, N is the number of observations.

The second step of the test of H_0 is to compare the calculated values of F to the standard table value, $F_{\alpha, p-1, N-p}$, for the desired level of confidence, α , with $p-1$

and $N-p$ degrees of freedom for upper and lower, respectively. If the calculated value of F exceeds the standard table value,

$F_{\alpha, p-1, N-p}$, then the null hypothesis is rejected at α level of significance and it is therefore inferred that the variation accounted for by the model (through the values of $\beta_i, i \neq 0$) is significant greater than the unexplained variation. In other words, the regression is significant if the calculated F -value is greater than the value obtained from standard table. It should be noted that this F -test is only valid for models for which there is no evidence of lack of fit.

Coefficient of determination and Absolute Average Deviation

An accompanying statistic to the F -test of Equation (A8) is the coefficient of determination:

$$R^2 = \frac{SS_R}{SS_T} \quad (A9)$$

The value of R^2 is a measure of the proportion of the total variation of the values of $y_{i, \text{exp}}$ about the mean $\bar{y}$ explained by the fitted model (Khuri and Cornel, 1987). Although this coefficient is a measure of how close the model fits the data, it cannot be used to judge model lack of fit because it does not take into account the number of degrees of freedom for model determination. In fact, a large value of R^2 does not necessarily mean that the regression model is a good one since addition of a variable to the model will always increase R^2 regardless of whether the additional variable is statistically significant or not. Thus, it is possible for models

that have large values of R^2 to yield poor predictions of new observations or estimates of the new response (Montgomery, 2005).

These types of errors can be eliminated by using the absolute average deviation (*AAD*) analysis, which is a direct method for describing the deviations (Bas and Boyaci, 2007). The *AAD* is calculated by the following equation:

$$AAD = \left\{ \left[\frac{\sum_{i=1}^N |y_{i,exp} - y_{i,cal}|}{y_{i,exp}} \right] / N \right\} \times 100 \quad (A10)$$

Where $y_{i,exp}$ and $y_{i,cal}$ are the experimental and calculated responses, respectively, and N is the number of experimental runs.

Evaluation of R^2 and *AAD* together is a better way of checking the adequacy of the model. R_2 must be close to 1.0 and *AAD* between the predicted and the observed data must be as small as possible. The acceptable values of R^2 and *AAD* values mean that the model equation defines the true behaviour of the system and it can be used for interpolation in the experimental domain, but, care must be taken about extrapolating beyond the region containing the original observations (Bas and Boyaci, 2007).

Test of hypothesis concerning individual parameters in the model

This test looks at the specific effects of the factors used in an experiment. The test of hypothesis concerning individual parameters in the model is performed by comparing the parameter estimates in the fitted model to their respective estimated standard errors (Khuri and Cornell, 1987). The test of the null hypothesis, $H_0: \beta_i=0$, is performed by calculating the value of the test statistic and comparing the value of t in Equation (A11) against a standard table value, t_α , from the t -tables:

$$t = \frac{\beta_i}{\text{est.s.e}(\beta_i)} \quad (\text{A11})$$

Where, β_i is the least squares estimate of coefficients, $\text{est.s.e.}(\beta_i)$ is the estimated standard errors of coefficient β_i . The estimated standard errors are the positive square roots of the products of the diagonal elements of $(X^T X)^{-1}$ and the mean sum of squares for residuals errors (MS_E) (Khuri and Cornell, 1987).

The choice of table value, t_α , depends on the alternative hypothesis, H_a , the level of significance, α , and the degree of freedom for t in Equation (A11). If the alternative hypothesis is, $H_a: \beta_i \neq 0$, the test is called two-sided test, and the value of $t_{\alpha/2}$ is taken from the column corresponding to in the standard table.

If, on the other hand, the alternative hypothesis is $H_a: \beta_i > 0$ or $H_a: \beta_i < 0$, the test is a one sided, and the value of t_α is taken from the column for t_α in the table.

If the absolute value of $t, |t|$, in Equation (A11) exceeds the standard table value, $t_{\alpha, N-p}$, then the null hypothesis, H_0 , is rejected in favour of the alternative, H_a

,(Khuri and Cornell, 1987); meaning the model coefficient is statistically significant.

Ridge Analysis

When the location of a stationary point is a saddle or is outside the experimental region, a search for the optimal value $\hat{y}$ is possible by the method of ridge analysis (Khuri and Cornell, 1987; Myers and Montgomery, 2002). In general, this method is used to finding the absolute maximum or minimum of $\hat{y}$ on concentric spheres of varying radii, R_i ($i=1,2, \dots$) which are centred at ($x_1, x_2, \dots, x_k$) = (0, 0, ..., 0) and are contained within the experimental region (Khuri and Cornell, 1987).

Given the fitted second order response surface model over the region of the k - coded variables ($x_1, x_2, \dots, x_k$) as,

$$\hat{y} = \beta_0 + x'b + x'\beta x \quad (A12)$$

If the coordinates of the variables that maximize $\hat{y}$ are restricted to the point lying on the boundary of a sphere of radius R , then the constraint becomes,

$$\sum_{i=1}^k x_i^2 = R^2 \quad (A13)$$

To maximize $\hat{y}$ in Equation (A12) subject to the constraint of Equation (A13), the function below is considered in accordance with the theorem by Kaplan (1952),

$$F = y - \mu(x'x - R^2)$$

or

$$F = \beta_0 + x'b + x'\beta x - \mu(x'x - R^2) \quad (A14)$$

Where: μ is the lagrangian multiplier and $x' = (x_1, x_2, \dots, x_k)$.

Differentiating Equation (A14) with respect to x_i , gives

$$\frac{\partial F}{\partial x} = b + 2\beta x - 2\mu x \quad (A15)$$

Equating Equation (A15) to zero and solving for x gives,

$$b - \mu I_k x = -\frac{b}{2}$$

or

$$x = -\frac{1}{2}(b - \mu I_k x) \quad (A16)$$

where I is the identity matrix.

The chosen values of μ can be inserted into Equation (A16), to solve for $x_1, x_2, \dots, x_k$, compute R from Equation (A13), and $\hat{y}$ from Equation (A12). However, it must be noted that the nature of the stationary point depends on the value of μ chosen (Khuri and Cornell, 1987).

The choice of the value of μ for generating a particular type of stationary point

There are two criteria for choosing the value of the lagrangian multiplier and they are given below (Myers *et al.*, 2009).

- If μ exceeds the largest eigenvalue of β in Equation Y, the solution x will results in an absolute maximum for $\hat{y}$ on $R = (x'x)^{\frac{1}{2}}$.
- If μ is smaller than the smallest eigenvalue of β , the solution x will results an absolute minimum for $\hat{y}$ on $R = (x'x)^{\frac{1}{2}}$.

APPENDIX B: RSM SAMPLE CALCULATIONS

Sample Calculation

The experimental data that was obtained from conducting this experiment is tabulated in Table B1.

Table B1: Experimental Data

Sample	Actual levels of variables			Experimental Ethanol Yield (Y _{exp})(mg/ml)
	Temperature (°C)	Sugars (g/500 ml)	Yeast (g/500ml)	
Factorial Points				
1	10	10	0.2	10.43
2	19	10	0.2	25.03
3	10	30	0.2	23.28
4	19	30	0.2	53.92
5	10	10	0.4	14.26
6	19	10	0.4	63.31
7	10	30	0.4	22.59
8	19	30	0.4	59.24
Axial Points				
9	7	20	0.3	28.89
10	22	20	0.3	48.39
11	14.5	wort	0.3	31.92
12	14.5	40	0.3	44.34
13	14.5	20	0.1	33.32
14	14.5	20	0.5	37.93
Centre Points				
15	14.5	20	0.3	29.22
16	14.5	20	0.3	33.31
17	14.5	20	0.3	27.36
18	14.5	20	0.3	32.59
19	14.5	20	0.3	36.93
20	14.5	20	0.3	26.82

The way in which sample 1 through to sample 20 were prepared is shown in the table. Ethanol concentrations for each of the samples were measured using the High Performance Liquid Chromatography (HPLC).

Having ethanol concentrations of samples 1 to 20, the average ethanol concentration was calculated as follows:

$$Y_{ave} = \frac{Y_{exp.1} + Y_{exp.2} + \dots + Y_{exp.19} + Y_{exp.20}}{20} \quad (B1)$$

$$Y_{ave} = 34.15 \text{ mg/ml}$$

where: Y_{ave} – Average ethanol concentration for sample 1 to sample 20

$Y_{exp.1}$ – Experimental ethanol concentration of sample 1

The Total Sum of Squares (SS_T) was then calculated as follows:

$$SS_T = \sum_{i=1}^{N=20} (y_{exp,i} - y_{ave})^2 \quad (B2)$$

$$= 3653.49$$

The degree of freedom of SS_T was calculated as follows:

$$df = N - 1 \quad (B3)$$

$$= 19$$

For the three variables under consideration, a second order polynomial regression model has been proposed by Myers *et al.* (2009) as:

$$y = \beta_0 + \sum_{i=1}^3 \beta_i X_i + \sum_{i=1}^3 \beta_{ii} X_i^2 + \sum_{i=1}^3 \sum_{j=i+1}^3 \beta_{ij} X_i X_j + \varepsilon \quad (B4)$$

where: y – Predicted response

β_0 – Coefficient of intercept

β_i – Coefficient of linear effect

β_{ii} – Coefficient of quadratic effect

β_{ij} – Coefficient of interaction effect

ε – Term that represents other sources of variability not accounted for by the response function

X_i & X_j – Coded independent variables

From here, MATLAB R2006a was used to calculate the coefficient values. The code for solving the coefficients as it appears in MATLAB R2007a is:

```
X=[1 -1 -1 -1 1 1 1 1 1 1;  
1 1 -1 -1 1 1 1 -1 -1 1;  
1 -1 1 -1 1 1 1 -1 1 -1;  
1 1 1 -1 1 1 1 1 -1 -1;  
1 -1 -1 1 1 1 1 1 -1 -1;  
1 1 -1 1 1 1 1 -1 1 -1;  
1 -1 1 1 1 1 1 -1 -1 1;  
1 1 1 1 1 1 1 1 1 1;  
1 -2^(3/4) 0 0 2^(3/2) 0 0 0 0 0;  
1 2^(3/4) 0 0 2^(3/2) 0 0 0 0 0;  
1 0 -2^(3/4) 0 0 2^(3/2) 0 0 0 0;  
1 0 2^(3/4) 0 0 2^(3/2) 0 0 0 0;  
1 0 0 -2^(3/4) 0 0 2^(3/2) 0 0 0;  
1 0 0 2^(3/4) 0 0 2^(3/2) 0 0 0;  
1 0 0 0 0 0 0 0 0 0;  
1 0 0 0 0 0 0 0 0 0;  
1 0 0 0 0 0 0 0 0 0;  
1 0 0 0 0 0 0 0 0 0;  
1 0 0 0 0 0 0 0 0 0];  
O=[1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1];  
I=eye(10);
```

Y=[10.42906;25.02525;23.27641;53.92276;14.2597;63.3093;22.594323;59.24323
;28.89464;48.39307;31.92476;44.34077;33.31592;37.92943;29.21793;33.31495;2
7.35822;32.59115;36.92509;26.8211];

N=20;p=10;n=6;

T=X'*X

U=X'*Y

V=inv(X'*X)

B=(X'*X)\(X'*Y)

Thus from Matlab R2007a, the values of the coefficients were found to be:

Table B2: Coefficient Values of the fitted model

β terms	Values
0	31.16
1	11.99
2	4.90
3	3.99
4	1.88
5	1.70
6	0.81
7	0.46
8	5.06
9	-4.68

The fitted model equation thus becomes:

$$y = 31.16 + 11.99x_1 + 4.90x_2 + 3.99x_3 + 1.88x_1^2 + 1.70x_2^2 + 0.81x_3^2 + 0.46x_1x_2 + 5.06x_1x_3 - 4.68x_2x_3$$

(B5)

Where: y – Ethanol concentration predicted by the fitted model

x – Coded values

Using this fitted model, predicted values of ethanol concentration for samples 1 to 20 were found to be as those tabulated in the table below and using their various corresponding coded values of x as seen also in the table below.

Table B3: Fitted Model ethanol concentration yield

Sample	Coded values of variables			Fitted Model Ethanol Yield (Y _{cal})(mg/ml)
	Temperatur e	Sugar s	Yeast	
Factorial Points				
1	-1	-1	-1	16.09
2	1	-1	-1	27.85
3	-1	1	-1	34.35
4	1	1	-1	47.92
5	-1	-1	1	22.13
6	1	-1	1	56.51
7	-1	1	1	21.65
8	1	1	1	57.85
Axial Points				
9	-1.682	0	0	16.31
10	1.682	0	0	56.64
11	0	-1.682	0	27.72
12	0	1.682	0	44.20
13	0	0	-1.682	26.74
14	0	0	1.682	40.16
Centre Points				
15	0	0	0	31.16
16	0	0	0	31.16
17	0	0	0	31.16
18	0	0	0	31.16
19	0	0	0	31.16
20	0	0	0	31.16

Having generated the fitted model, the Sum of Squares of Regression was calculated as:

$$SS_R = \sum_{i=1}^{N=20} (y_{cal,i} - y_{ave})^2 \quad (B6)$$

$$= 3029.18$$

The degree of freedom of SS_R was calculated as follows:

$$df = p - 1$$

$$= 9$$

where p – number of parameters in the model, which is equals to 10 for this model

The Sum of Squares of Errors (SS_E) was calculated as:

$$SS_E = \sum_{i=1}^{N=20} (y_{\text{exp},i} - y_{\text{cal},i})^2 \quad (\text{B7})$$

$$= 678.99$$

The degree of freedom of SS_E was calculated as follows:

$$df = N - p \quad (\text{B8})$$

$$= 10$$

The centre points average ethanol concentration yield was calculated as follows:

$$y_{\text{ave,centre}} = \frac{y_{\text{exp},15} + y_{\text{exp},16} + \dots + y_{\text{exp},20}}{6} \quad (\text{B9})$$

$$= 31.04 \text{ mg/ml}$$

Test for Lack of Fit for the fitted model was then carried out as follows:

Having calculated the average value of ethanol concentration for the experimental centre point (sample 15 to sample 20), Sum of Squares due to pure error (SS_{PE}) was then calculated as;

$$SS_{PE} = \sum_{\text{centre-point}} (y_{\text{exp,centre-point}} - y_{\text{ave,centre}})^2 \quad (\text{B10})$$

where: $y_{\text{exp,centre}}$ – observations of the centre points

$Y_{\text{ave,centre}}$ – average of observations of the n_c runs at the centre points

The degree of freedom of SS_{PE} was calculated as follows:

$$df = n_c - 1$$

$$= 5$$

where: n_c – is the number of observations at the centre points, which is equals to 6

By subtraction, the sum of squares due to lack of fit (SS_{LOF}) is then calculated as:

$$SS_{LOF} = SS_E - SS_{PE} \quad (\text{B11})$$

$$= 602.10$$

The degree of freedom of SS_{LOF} was calculated as follows:

$$df = N - p - n_c + 1 \quad (\text{B12})$$

$$= 5$$

Using the SS_{LOF} , the Fisher's variance ratio test (F-test) to test for lack of fit was then calculated as:

$$F = \frac{\left(\frac{SS_{LOF}}{N - p - n_c + 1} \right)}{\left(\frac{SS_{PE}}{n_c - 1} \right)} \quad (B13)$$

$$= 7.83$$

This value was then compared with the standard table value of F ratio for desired level of confidence of $\alpha = 0.01$, which was found to be 10.97 (Khuri and Cornell, 1987). The F-test value calculated as seen is less than that of the standard table value of 10.97, thus this proves that the fitted model has zero lack of fit.

The F-test for significance of regression

The test for significance of regression is a test to determine whether a linear relationship exists between the response variable, y , and the regression coefficients, β_i (Montgomery, 2005). The usual test of the significance of the fitted regression equation is a test of the null hypothesis, H_0 : all values of β_i (excluding β_0) are zero. The alternative hypothesis is, H_a : at least one value of β_i (excluding β_0) is not zero (Simate, 2009; Myers and Montgomery, 2002). Assuming normality of the errors, the test of H_0 involves first calculating the value of the F-statistic as:

$$F = \text{Mean Square of Regression} / \text{Mean Square residual} \quad (B14)$$

$$= \frac{\left(\frac{SS_R}{p-1} \right)}{\left(\frac{SS_E}{N-p} \right)}$$

$$= 4.96$$

This value was then compared with the standard table value of F for the desired level of confidence of $\alpha = 0.01$, which was found to be 4.94 (Khuri and Cornell, 1987). The F-test value calculated as seen is more than that of the standard table value of 4.94, thus this proves that the regression is significant at a confidence level of 99%.

The value of R^2 which is a measure of the proportion of the total variation of the values of $y_{i,exp}$ about the mean y_{ave} explained by the fitted (Khuri and Cornell, 1987), was calculated as;

$$R^2 = \frac{SS_R}{SS_T} \quad (B15)$$

$$= 0.8291$$

This shows that 82.91% of ethanol production can be explained by this model (Khuri and Cornell, 1987).

Test of hypothesis concerning individual parameters in the model.

This test uses the t-test method, whereby t is calculated as;

$$t = \frac{\beta_i}{est.s.e(\beta_i)} \quad (B16)$$

Where β_i is the least squares estimate of coefficients, est.s.e. (β_i) is the estimated standard errors of coefficient β_i .

The estimated standard errors are the positive square roots of the products of the diagonal elements of $(X'X)^{-1}$ and the mean sum of squares for residuals (errors), MS_E (Khuri and Cornell, 1987).

Thus MS_E was calculated as;

$$MS_E = \frac{SS_E}{N - p} \quad (B17)$$

$$= 67.90$$

Having calculated the MS_E , $est.s.e(\beta)$ was calculated as;

$$est.s.e(\beta_0) = \sqrt{MS_E \times Diagonal_{Matrix}} \quad (B18)$$

$$= 3.36$$

The same thing was done for all the coefficients and is tabulated in the table below. Then the t value was then calculated according to the equation above.

Table B4: t-Test values

β	Coefficient Value (β_i)	Diagonal Matrix	est.s.e.(β)	t-Test
0	31.1623	0.166	3.36	9.28
1	11.9891	0.073	2.23	5.39
2	4.8982	0.073	2.23	2.20
3	3.9916	0.073	2.23	1.79
4	1.8768	0.069	2.16	0.87
5	1.6961	0.069	2.16	0.78
6	0.8087	0.069	2.16	0.37
7	0.4562	0.125	2.91	0.16
8	5.657	0.125	2.91	1.94
9	-4.6845	0.125	2.91	-1.61

The t values were then compared to that of the standard table value obtained from (Khuri and Cornell, 1987) as 2.76. According to Khuri and Cornell (1987), the t values calculated above should be greater than that obtained from the standard table of 2.76 for the corresponding term to be significant. Thus it can be deduced from the table above that only the constant term (β_0) and the first linear term (β_1) which represents temperature are the once with a value greater than that obtained from the standard table, meaning that they are the once with the most significant effect at 99% confidence level, and the other terms are negligible but have been proofed to also have some effect as it was shown from the graphs in the results section.

Program for determining the regression coefficients in a fitted model

```
X=[1 -1 -1 -1 1 1 1 1 1 1;  
  
1 1 -1 -1 1 1 1 -1 -1 1;  
1 -1 1 -1 1 1 1 -1 1 -1;  
1 1 1 -1 1 1 1 1 -1 -1;  
1 -1 -1 1 1 1 1 1 -1 -1;  
1 1 -1 1 1 1 1 -1 1 -1;  
1 -1 1 1 1 1 1 -1 -1 1;  
1 1 1 1 1 1 1 1 1 1;  
1 -2^(3/4) 0 0 2^(3/2) 0 0 0 0 0;  
1 2^(3/4) 0 0 2^(3/2) 0 0 0 0 0;  
1 0 -2^(3/4) 0 0 2^(3/2) 0 0 0 0;  
1 0 2^(3/4) 0 0 2^(3/2) 0 0 0 0;  
1 0 0 -2^(3/4) 0 0 2^(3/2) 0 0 0;  
1 0 0 2^(3/4) 0 0 2^(3/2) 0 0 0;  
1 0 0 0 0 0 0 0 0 0;  
1 0 0 0 0 0 0 0 0 0;  
1 0 0 0 0 0 0 0 0 0;  
1 0 0 0 0 0 0 0 0 0;  
1 0 0 0 0 0 0 0 0 0;  
1 0 0 0 0 0 0 0 0 0];  
O=[1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1];  
I=eye(10);  
Y=[10.42906;25.02525;23.27641;53.92276;14.2597;63.3093;22.594323;59.24323  
;28.89464;48.39307;31.92476;44.34077;33.31592;37.92943;29.21793;33.31495;2  
7.35822;32.59115;36.92509;26.8211];  
N=20;p=10;n=6;  
T=X'*X
```

```
U=X'*Y
```

```
V=inv(X'*X)
```

```
B=(X'*X)\(X'*Y)
```

Programme for plotting response surfaces and contours

Response surface and contour at constant Sugar concentration

RESPONSE SURFACE

```
[x,y]=meshgrid (-1.682:1:1.682,-1.682:1:1.682);  
z=31.1623+11.9891*x+3.9916*y+1.8768*(x.^2)+0.8087*(y.^2)+5.0570*x*y;  
mesh(x,y,z)  
xlabel('Temperature'),ylabel('Yeast Conc'),zlabel('Ethanol Conc')
```

CONTOUR PLOT

```
[x,y]=meshgrid(-1.682:1:1.682,-1.682:1:1.682);  
z=31.1623+11.9891*x+3.9916*y+1.8768*(x.^2)+0.8087*(y.^2)+5.0570*x*y;  
[C,h]=contour(x,y,z);  
xlabel('Temperature'),ylabel('Yeast Conc'),zlabel('Ethanol Conc')  
set(h,'ShowText','on','TextStep',get(h,'LevelStep')*2)  
colormap cool
```

Response surface and contour at constant yeast concentration

RESPONSE SURFACE

```
[x,y]=meshgrid(-1.682:1:1.682,-1.682:1:1.682);  
z=31.1623+11.9891*x+4.8982*y+1.8768*(x.^2)+1.6961*(y.^2)+0.4562*x*y;  
mesh(x,y,z)  
xlabel('Temperature'),ylabel('Sugar Conc'),zlabel('Ethanol Conc')
```

CONTOUR PLOT

```
[x,y]=meshgrid(-1.682:1:1.682,-1.682:1:1.682);  
z=31.1623+11.9891*x+4.8982*y+1.8768*(x.^2)+1.6961*(y.^2)+0.4562*x*y;  
[C,h]=contour(x,y,z);  
xlabel('Temperature'),ylabel('Sugar Conc'),zlabel('Ethanol Conc')  
set(h,'ShowText','on','TextStep',get(h,'LevelStep')*2)  
colormap cool
```

Response surface and contour at constant Temperature

RESPONSE SURFACE

```
[x,y]=meshgrid(-1.682:1:1.682,-1.682:1:1.682);  
z=31.1623+4.8982*x+3.9916*y+1.6961*(x.^2)+0.8087*(y.^2)-4.6845*x*y;  
mesh(x,y,z)  
xlabel('Sugar Conc'),ylabel('Yeast Conc'),zlabel('Ethanol Conc')
```

CONTOUR PLOT

```
[x,y]=meshgrid(-1.682:1:1.682,-1.682:1:1.682);  
z=31.1623+4.8982*x+3.9916*y+1.6961*(x.^2)+0.8087*(y.^2)-4.6845*x*y;  
[C,h]=contour(x,y,z);  
xlabel('Sugar Conc'),ylabel('Yeast Conc'),zlabel('Ethanol Conc')  
set(h,'ShowText','on','TextStep',get(h,'LevelStep')*2)  
colormap cool
```

APPENDIX C: HIGH PERFORMANCE LIQUID CHROMATOGRAPHY

DATA

HPLC Sugar analysis

Wort Sample

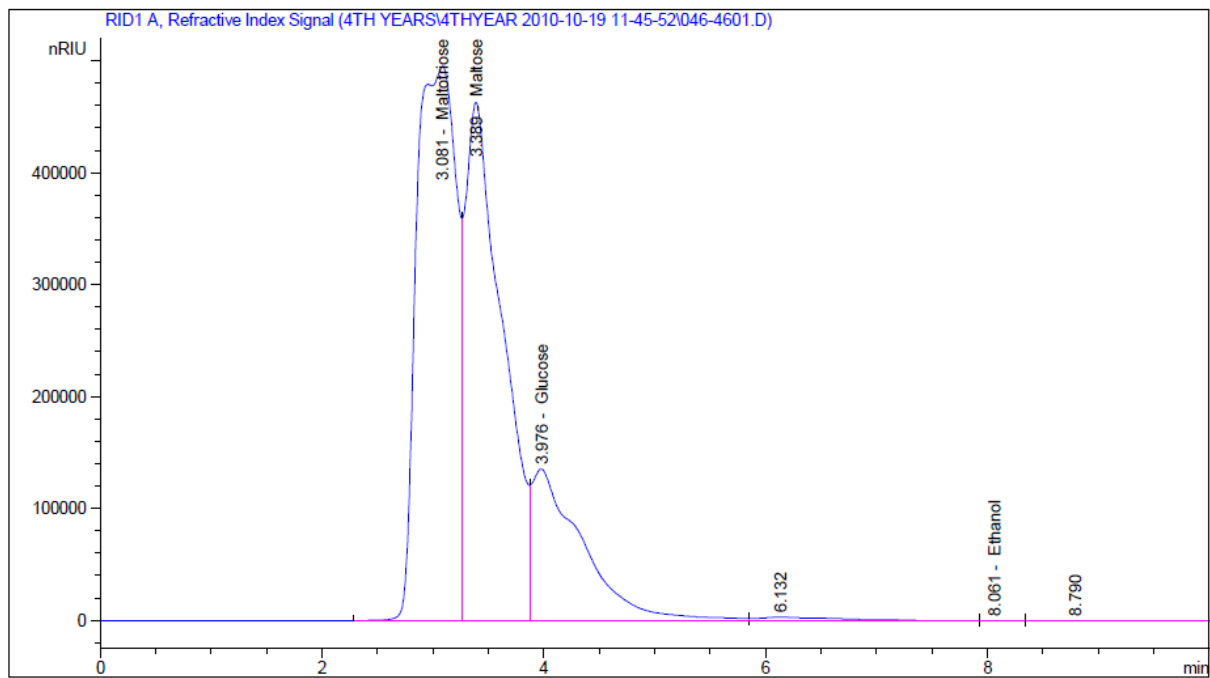


Figure C1: Retention Time for Wort Sample

Table C1: Sugar Concentrations for wort sample

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
3.081	BV	1.22521e7	1.53886e-5	188.54322		Maltotriose
3.389	VV	1.09775e7	9.93151e-6	109.02360		Maltose
3.976	VV	4.29112e6	9.14063e-6	39.22351		Glucose
4.325		-	-	-		Fructose
5.562		-	-	-		Glycerol
6.132	VV	1.50306e5	0.00000	0.00000	?	
8.061	VV	6046.36523	3.41902e-5	2.06726e-1		Ethanol
8.790	VBA	2.48933e4	0.00000	0.00000	?	
Totals :				336.99706		

Only Maltose and Glucose were considered for the sugars, thus their concentrations were summed up to be 148 mg/ml of sugar.

Wort plus 10g of sugar analysis

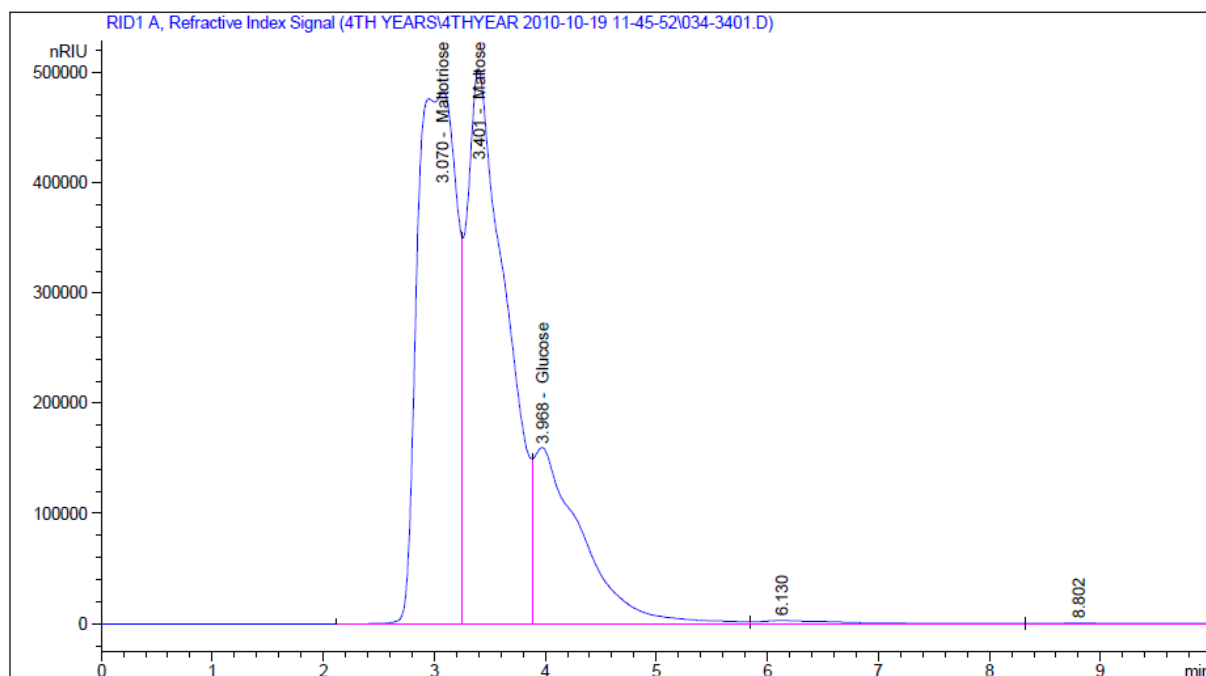


Figure C2: Retention Time for Wort plus 10g of sugar Sample

Table C2: Sugar Concentrations for wort plus 10g of sugar sample

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
3.070	VV	1.19067e7	1.53886e-5	183.22810		Maltotriose
3.401	VV	1.25129e7	9.93693e-6	124.33984		Maltose
3.968	VV	4.80913e6	9.17156e-6	44.10728		Glucose
4.325		-	-	-		Fructose
5.562		-	-	-		Glycerol
6.130	VV	1.62986e5	0.00000	0.00000		?
8.120		-	-	-		Ethanol
8.802	VBA	3.37469e4	0.00000	0.00000		?
Totals :				351.67522		

Sugar conc = 168 mg/ml

Wort plus 20g of sugar analysis

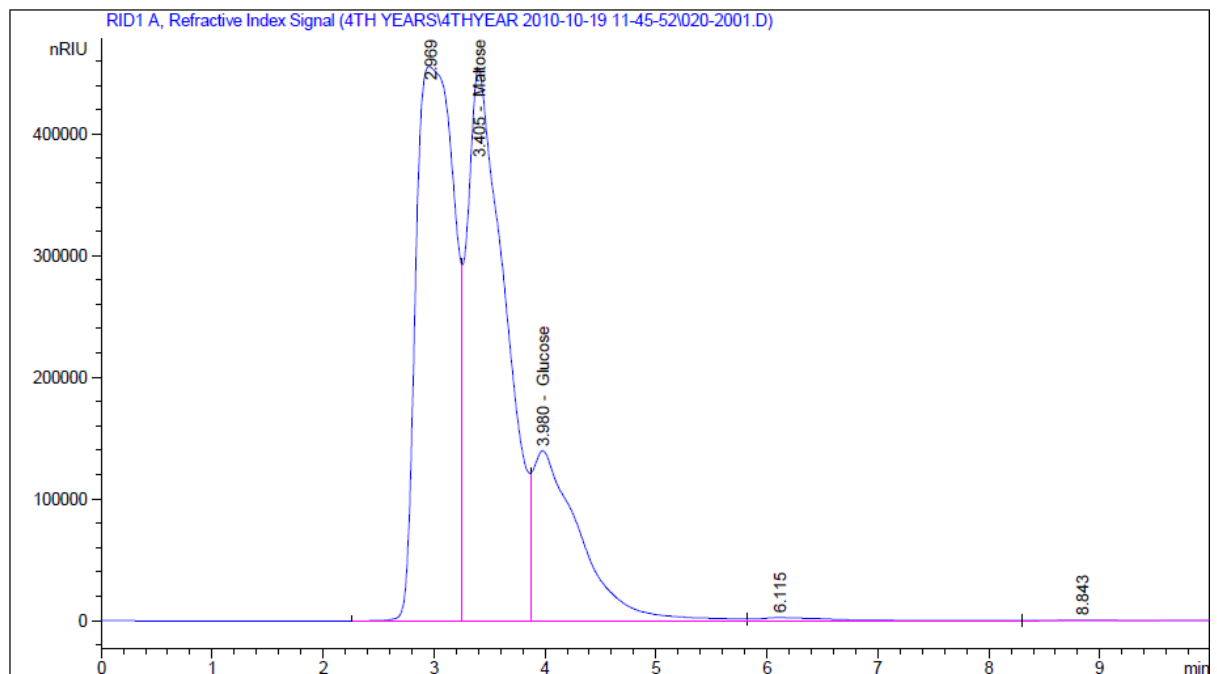


Figure C3: Retention Time for Wort plus 20g of sugar Sample

Table C3: Sugar Concentrations for wort plus 20g of sugar sample

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
2.969	BV	1.09439e7	0.00000	0.00000	?	
3.080		-	-	-		Maltotriose
3.405	VV	1.09422e7	9.93137e-6	108.67146		Maltose
3.980	VV	4.17565e6	9.13269e-6	38.13487		Glucose
4.325		-	-	-		Fructose
5.562		-	-	-		Glycerol
6.115	VV	1.32994e5	0.00000	0.00000	?	
8.120		-	-	-		Ethanol
8.843	VBA	3.14165e4	0.00000	0.00000	?	
Totals :				146.80633		

Sugar conc = 147 mg/ml

Wort plus 30g of sugar analysis

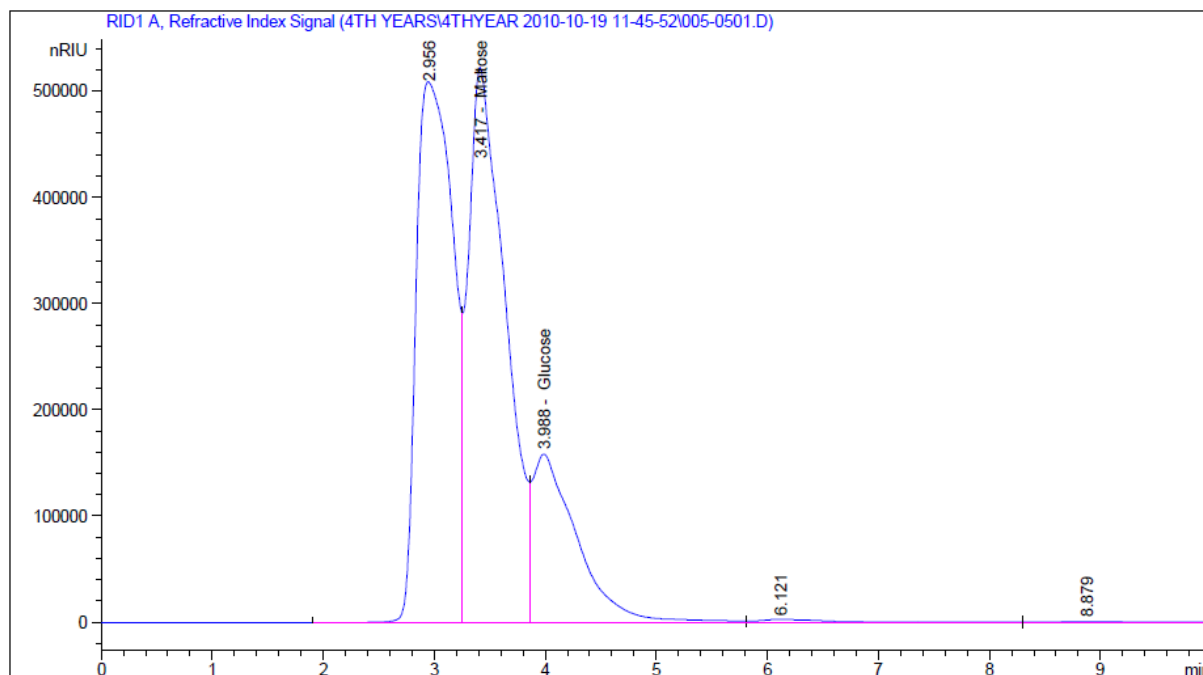


Figure C4: Retention Time for Wort plus 30g of sugar Sample

Table C4: Sugar Concentrations for wort plus 30g of sugar sample

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
2.956	VV	1.16493e7	0.00000	0.00000	?	
3.080		-	-	-		Maltotriose
3.417	VV	1.22529e7	9.93610e-6	121.74576		Maltose
3.988	VV	4.41459e6	9.14866e-6	40.38761		Glucose
4.325		-	-	-		Fructose
5.562		-	-	-		Glycerol
6.121	VV	1.28825e5	0.00000	0.00000	?	
8.120		-	-	-		Ethanol
8.879	VBA	4.02598e4	0.00000	0.00000	?	
Totals :				162.13337		

Sugar conc = 162 mg/ml

Wort plus 40g of sugar analysis

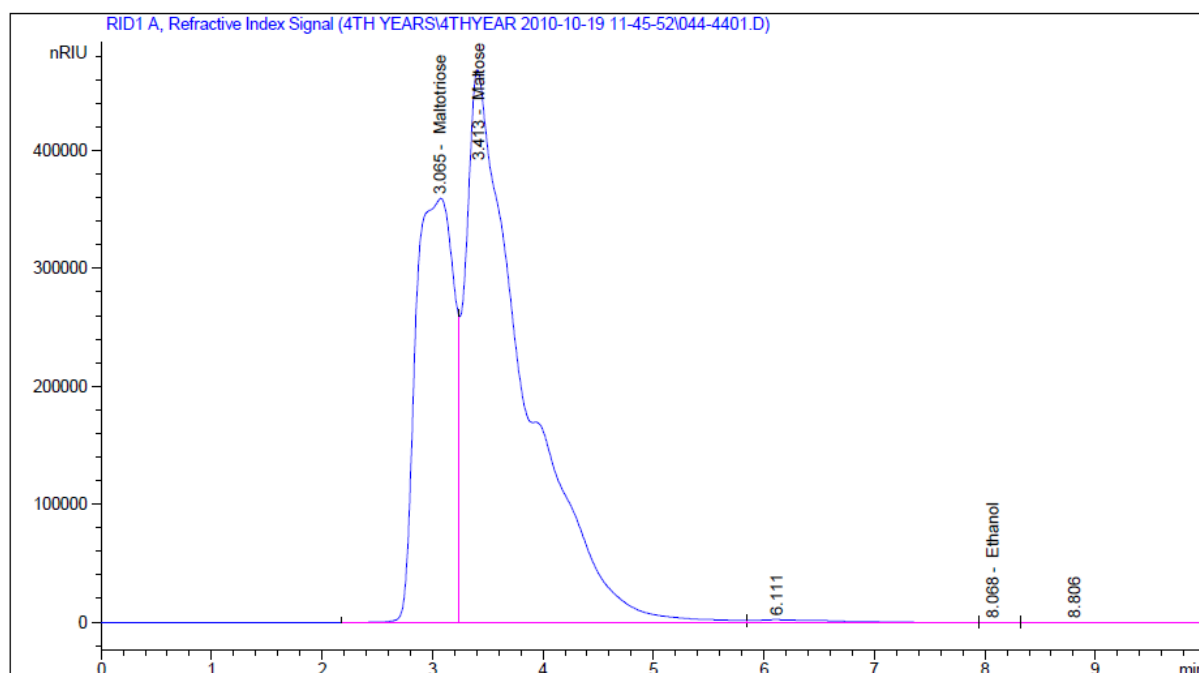


Figure C5: Retention Time for Wort plus 40g of sugar Sample

Table C5: Sugar Concentrations for wort plus 40g of sugar sample

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
3.065	VV	8.56703e6	1.53886e-5	131.83470		Maltotriose
3.413	VV	1.73647e7	9.94774e-6	172.73996		Maltose
3.984		-	-	-		Glucose
4.325		-	-	-		Fructose
5.562		-	-	-		Glycerol
6.111	VV	1.14001e5	0.00000	0.00000		?
8.068	VV	2958.26343	4.65259e-5	1.37636e-1		Ethanol
8.806	VBA	1.77761e4	0.00000	0.00000		?
Totals :				304.71230		

Sugar conc = 173 mg/ml

Ethanol analysis

Sample 1: Temperature = 10 °C; Sugar = 10g; Yeast = 0.2g

Table C6: Ethanol Concentrations for sample 1

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
2.960	VV	1.15058e7	0.00000	0.00000		?
3.080		-	-	-		Maltotriose
3.393	VV	7.78738e6	9.91343e-6	77.19961		Maltose
3.983	VV	1.05183e6	8.25617e-6	8.68406		Glucose
4.266	VV	1.32185e6	9.05630e-6	11.97103		Fructose
5.581	VV	1.18783e5	1.07899e-5	1.28165		Glycerol
6.084	VV	1.20699e5	0.00000	0.00000		?
6.997	VV	3.58173e4	0.00000	0.00000		?
8.229	VBA	4.62949e5	2.25274e-5	10.42906		Ethanol
Totals :				109.56541		

Sample 2: Temperature = 19°C; Sugar = wort +10g sugar; Yeast = 0.2g

Table C7: Ethanol Concentrations for sample 2

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
2.924	BV	7.13448e6	0.00000	0.00000	?	
3.080		-	-	-		Maltotriose
3.379		-	-	-		Maltose
3.976	VV	1.28186e5	0.00000	0.00000		Glucose
4.325		-	-	-		Fructose
5.163	VV	2.99813e4	0.00000	0.00000	?	
5.590	VV	1.86408e5	1.11115e-5	2.07128		Glycerol
7.371	VV	5489.67480	0.00000	0.00000	?	
8.197	VBA	1.11535e6	2.24372e-5	25.02525		Ethanol
Totals :				27.09652		

Sample 3: Temp = 10 °C; Sugar = wort +30g sugar; Yeast = 0.2g

Table C8: Ethanol Concentrations for sample 3

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
3.026	VV	8.92138e6	1.53886e-5	137.28760		Maltotriose
3.389	VV	8.55375e6	9.91900e-6	84.84471		Maltose
3.984		-	-	-		Glucose
4.323	VV	2.28532e6	9.07127e-6	20.73071		Fructose
5.586	VV	3.62202e5	1.13857e-5	4.12393		Glycerol
7.366	VV	2.11576e4	0.00000	0.00000	?	
8.211	VBA	1.06696e6	2.24401e-5	23.94264		Ethanol
Totals :				270.92959		

Sample 4: Temp = 19 °C; Sugar = wort +30g sugar; Yeast = 0.2g

Table C9: Ethanol Concentrations for sample 4

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
2.951	BV	1.28981e7	0.00000	0.00000	?	
3.080		-	-	-		Maltotriose
3.379		-	-	-		Maltose
3.984		-	-	-		Glucose
4.325		-	-	-		Fructose
4.950	VV	2.53527e4	0.00000	0.00000	?	
5.184	VV	3.51960e4	0.00000	0.00000	?	
5.600	VV	4.28344e5	1.14306e-5	4.89623		Glycerol
7.397	VV	3.06616e4	0.00000	0.00000	?	
8.213	VBA	2.40697e6	2.24028e-5	53.92276		Ethanol

Sample 5: Temp = 10 °C; Sugar = wort +10g sugar; Yeast = 0.4g

Table C10: Ethanol Concentrations for sample 5

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
2.983	BV	8.61286e6	0.00000	0.00000	?	
3.080		-	-	-		Maltotriose
3.379	VV	6.52980e6	9.90145e-6	64.65452		Maltose
3.984		-	-	-		Glucose
4.312	VV	1.29675e6	9.05561e-6	11.74287		Fructose
5.584	VV	1.23113e5	1.08211e-5	1.33222		Glycerol
6.011	VV	1.21709e5	0.00000	0.00000	?	
6.975	VV	3.06371e4	0.00000	0.00000	?	
8.212	VBA	6.34160e5	2.24858e-5	14.25957		Ethanol
Totals :				91.98917		

Sample 6: Temp = 19 °C; Sugar = wort +10g sugar; Yeast = 0.4g

Table C11: Ethanol Concentrations for sample 6

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
2.957	BV	1.70585e7	0.00000	0.00000	?	
3.080		-	-	-		Maltotriose
3.379		-	-	-		Maltose
3.984		-	-	-		Glucose
4.325		-	-	-		Fructose
5.180	VV	4.40765e4	0.00000	0.00000	?	
5.596	VV	5.52038e5	1.14857e-5	6.34054		Glycerol
7.380	VV	2.84308e4	0.00000	0.00000	?	
8.197	VBA	2.82651e6	2.23984e-5	63.30930		Ethanol
Totals :				69.64984		

Sample 7: Temp = 10 °C; Sugar = wort +30g sugar; Yeast = 0.4g

Table C12: Ethanol Concentrations for sample 7

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
2.969	BV	1.02013e7	0.00000	0.00000	?	
3.080		-	-	-		Maltotriose
3.385	VV	8.88358e6	9.92111e-6	88.13494		Maltose
3.984		-	-	-		Glucose
4.324	VV	2.38779e6	9.07215e-6	21.66240		Fructose
5.584	VV	3.70675e5	1.13924e-5	4.22287		Glycerol
7.095	VV	4.38908e4	0.00000	0.00000	?	
8.211	VBA	9.93891e5	2.24450e-5	22.30787		Ethanol
Totals :				136.32807		

Sample 8: Temp = 19 °C; Sugar = wort +30g sugar; Yeast = 0.4g

Table C13: Ethanol Concentrations for sample 8

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
2.950	BV	1.16849e7	0.00000	0.00000	?	
3.080		-	-	-		Maltotriose
3.379		-	-	-		Maltose
3.984		-	-	-		Glucose
4.325		-	-	-		Fructose
4.960	VV	3.23778e4	0.00000	0.00000	?	
5.177	VV	3.40922e4	0.00000	0.00000	?	
5.586	VV	4.45390e5	1.14400e-5	5.09527		Glycerol
7.380	VV	2.41444e4	0.00000	0.00000	?	
8.190	VBA	2.64477e6	2.24001e-5	59.24323		Ethanol

Sample 9: Temp = 7 °C; Sugar = wort +20g sugar; Yeast = 0.3g

Table C14: Ethanol Concentrations for sample 9

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
2.940	BV	9.78426e6	0.00000	0.00000	?	
3.080		-	-	-		Maltotriose
3.410	VV	8.58621e6	9.91922e-6	85.16848		Maltose
3.998	VV	4.79172e6	9.17063e-6	43.94311		Glucose
4.325		-	-	-		Fructose
5.562	VV	5.61239e4	9.80010e-6	5.50020e-1		Glycerol
6.105	VV	1.08911e5	0.00000	0.00000	?	
8.312	VBA	1.28830e6	2.24286e-5	28.89464		Ethanol
Totals :				158.55625		

Sample 10: Temp = 22 °C; Sugar = wort +20g sugar; Yeast = 0.3g

Table C15: Ethanol Concentrations for sample 10

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
2.948	VV	1.15664e7	0.00000	0.00000	?	
3.080		-	-	-		Maltotriose
3.379		-	-	-		Maltose
3.984		-	-	-		Glucose
4.325		-	-	-		Fructose
5.175	VV	8.78353e4	0.00000	0.00000	?	
5.602	VV	3.78978e5	1.13986e-5	4.31982		Glycerol
7.014	VV	2.74785e4	0.00000	0.00000	?	
7.304	VV	2.55806e4	0.00000	0.00000	?	
8.221	VBA	2.15981e6	2.24062e-5	48.39307		Ethanol

Sample 11: Temp = 14.5 °C; Sugar = wort ; Yeast = 0.3g

Table C17: Ethanol Concentrations for sample 11

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
2.950	VV	1.39266e7	0.00000	0.00000	?	
3.080		-	-	-		Maltotriose
3.379		-	-	-		Maltose
3.984		-	-	-		Glucose
4.325		-	-	-		Fructose
4.949	VV	8.25497e4	0.00000	0.00000	?	
5.590	VV	3.18885e5	1.13462e-5	3.61814		Glycerol
7.022	VV	1.42323e5	0.00000	0.00000	?	
8.215	VBA	1.42373e6	2.24233e-5	31.92476		Ethanol
Totals :				35.54290		

Sample 12: Temp = 14.5 °C; Sugar = wort +40g sugar; Yeast = 0.3g

Table C18: Ethanol Concentrations for sample 12

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
2.948	BV	8.95622e6	0.00000	0.00000	?	
3.080		-	-	-		Maltotriose
3.394	VV	6.19222e6	9.89740e-6	61.28693		Maltose
3.984		-	-	-		Glucose
4.318	VV	4.46535e5	8.98669e-6	4.01287		Fructose
5.599	VV	3.98817e5	1.14124e-5	4.55146		Glycerol
6.984	VV	1.27441e4	0.00000	0.00000	?	
7.377	VV	2.63142e4	0.00000	0.00000	?	
8.218	VBA	1.97868e6	2.24092e-5	44.34077		Ethanol
Totals :				114.19204		

Sample 13: Temp = 14.5 °C; Sugar = wort +20g sugar; Yeast = 0.1g

Table C19: Ethanol Concentrations for sample 13

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
2.949	BV	9.98712e6	0.00000	0.00000	?	
3.080		-	-	-		Maltotriose
3.396	VV	5.98659e6	9.89472e-6	59.23564		Maltose
3.984		-	-	-		Glucose
4.322	VV	3.32160e5	8.95050e-6	2.97300		Fructose
5.599	VV	3.25082e5	1.13525e-5	3.69049		Glycerol
7.008	VV	4.29469e4	0.00000	0.00000	?	
8.218	VBA	1.48591e6	2.24212e-5	33.31592		Ethanol
Totals :				99.21505		

Sample 14: Temp = 14.5 °C; Sugar = wort +20g sugar; Yeast = 0.5g

Table C 20: Ethanol Concentrations for sample 14

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
2.955	VV	1.11460e7	0.00000	0.00000	?	
3.080		-	-	-		Maltotriose
3.376	VV	5.36166e6	9.88528e-6	53.00151		Maltose
3.984		-	-	-		Glucose
4.325		-	-	-		Fructose
4.998	VV	9.46943e4	0.00000	0.00000	?	
5.598	VV	3.74721e5	1.13954e-5	4.27010		Glycerol
7.010	VV	6.48631e4	0.00000	0.00000	?	
8.222	VBA	1.69212e6	2.24153e-5	37.92943		Ethanol
Totals :				95.20104		

Sample 15: Temp = 14.5 °C; Sugar = wort +20g sugar; Yeast = 0.3g

Table C 21: Ethanol Concentrations for sample 15

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
2.939	BV	8.61959e6	0.00000	0.00000	?	
3.080		-	-	-		Maltotriose
3.392	VV	4.31287e6	9.86331e-6	42.53919		Maltose
3.984		-	-	-		Glucose
4.295	VV	1.25291e5	8.71719e-6	1.09218		Fructose
4.997	VV	7.05308e4	0.00000	0.00000	?	
5.597	VV	2.86927e5	1.13094e-5	3.24498		Glycerol
7.020	VV	7.43479e4	0.00000	0.00000	?	
8.211	VBA	1.30275e6	2.24279e-5	29.21793		Ethanol
Totals :				76.09429		

Sample 16: Temp = 14.5 °C; Sugar = wort +20g sugar; Yeast = 0.3g

Table C 22: Ethanol Concentrations for sample 16

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
2.945	VV	9.77234e6	0.00000	0.00000	?	
3.080		-	-	-		Maltotriose
3.390	VV	5.79428e6	9.89203e-6	57.31722		Maltose
3.984		-	-	-		Glucose
4.325		-	-	-		Fructose
5.166	VV	3.41106e4	0.00000	0.00000	?	
5.596	VV	3.33379e5	1.13606e-5	3.78738		Glycerol
6.946	VV	8976.69434	0.00000	0.00000	?	
7.365	VV	1.54348e4	0.00000	0.00000	?	
8.217	VBA	1.48587e6	2.24212e-5	33.31495		Ethanol

Sample 17: Temp = 14.5 °C; Sugar = wort +20g sugar; Yeast = 0.3g

Table C 23: Ethanol Concentrations for sample 17

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
2.952	VV	8.38018e6	0.00000	0.00000	?	
3.080		-	-	-		Maltotriose
3.376	VV	4.22719e6	9.86103e-6	41.68446		Maltose
3.984		-	-	-		Glucose
4.325		-	-	-		Fructose
4.966	VV	7.09382e4	0.00000	0.00000	?	
5.595	VV	2.73100e5	1.12909e-5	3.08354		Glycerol
7.026	VV	1.04214e5	0.00000	0.00000	?	
8.223	VBA	1.21962e6	2.24317e-5	27.35822		Ethanol
Totals :				72.12622		

Sample 18: Temp = 14.5 °C; Sugar = wort +20g sugar; Yeast = 0.3g

Table C 24: Ethanol Concentrations for sample 18

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
2.955	VV	9.26143e6	0.00000	0.00000	?	
3.080		-	-	-		Maltotriose
3.381	VV	6.11814e6	9.89646e-6	60.54788		Maltose
3.984		-	-	-		Glucose
4.325		-	-	-		Fructose
5.595	VV	3.43418e5	1.13698e-5	3.90460		Glycerol
7.379	VV	1.72334e4	0.00000	0.00000	?	
8.219	VBA	1.45352e6	2.24223e-5	32.59115		Ethanol
Totals :				97.04363		

Sample 19: Temp = 14.5 °C; Sugar = wort +20g sugar; Yeast = 0.3g

Table C 25: Ethanol Concentrations for sample 19

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
2.956	VV	1.02145e7	0.00000	0.00000	?	
3.080		-	-	-		Maltotriose
3.376	VV	5.88847e6	9.89337e-6	58.25684		Maltose
3.984		-	-	-		Glucose
4.325		-	-	-		Fructose
4.972	VV	8.68994e4	0.00000	0.00000	?	
5.595	VV	3.76121e5	1.13965e-5	4.28645		Glycerol
7.014	VV	9.68447e4	0.00000	0.00000	?	
8.215	VBA	1.64723e6	2.24165e-5	36.92509		Ethanol
Totals :				99.46839		

Sample 20: Temp = 14.5 °C; Sugar = wort +20g sugar; Yeast = 0.3g

Table C 26: Ethanol Concentrations for sample 20

RetTime [min]	Type	Area [nRIU*s]	Amt/Area	Amount [mg/ml]	Grp	Name
2.965	BV	7.61935e6	0.00000	0.00000	?	
3.080		-	-	-		Maltotriose
3.370	VV	5.66525e6	9.89013e-6	56.03003		Maltose
3.984		-	-	-		Glucose
4.325		-	-	-		Fructose
5.583	VV	2.80943e5	1.13016e-5	3.17511		Glycerol
6.950	VV	1.69705e4	0.00000	0.00000	?	
7.320	VV	1.56020e4	0.00000	0.00000	?	
8.199	VBA	1.19562e6	2.24329e-5	26.82110		Ethanol
Totals :				86.02624		

APPENDIX D: DATA FOR ARTIFICIAL NEURAL NETWORKS

Table D1: Default training parameters for the Levenberg-Marquardt training function

Training parameter	Value assigned
epochs	100
goal	5
max_fail	5
mem_reduc	1
min_grad	1.0000e-10
mu	1.0000e-03
mu_dec	0.1000
mu_inc	10
mu_max	1.0000e+10
show	25

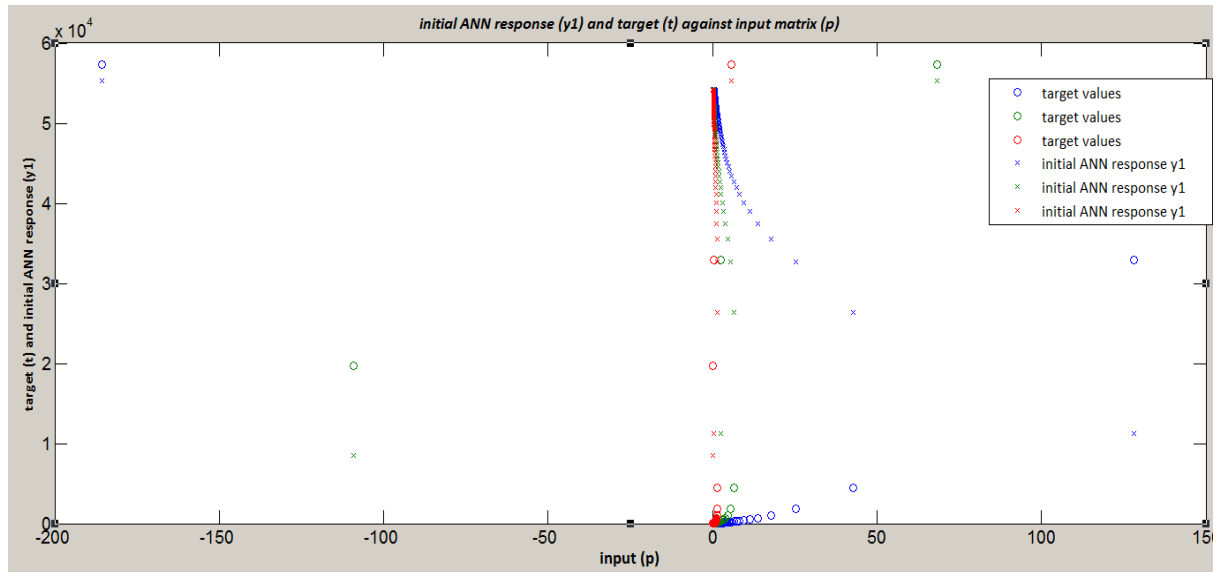


Figure D1 The ANN's initial response ('x') to the input matrix as compared to the target values ('o').

From Figure D1 above, it can be seen that the network's initial response to the input vector, p is slightly lower than the target (desired) values. In fact, they are almost symmetrical about the x-axis which implies that initially, more training was required to shift the y values upwards to reach more positive values hence approximating the target vector.

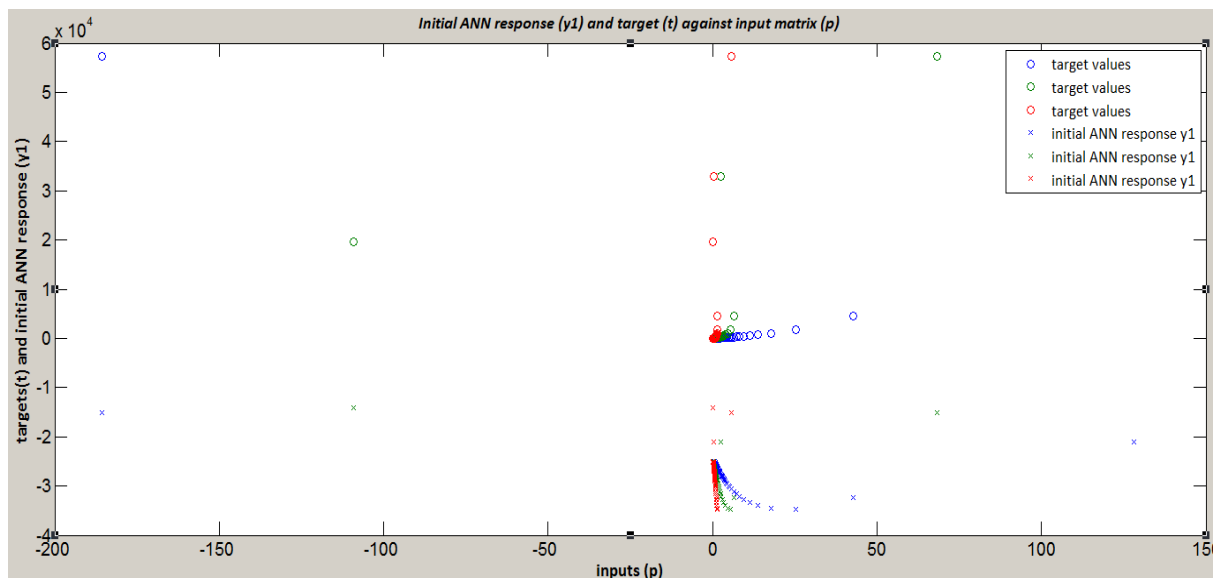


Figure D2: The ANN's initial response ('x') to the input matrix as compared to the target values ('o').

Figure D.2 above shows the network's initial response to the input versus the target once the network has been well trained. It can be deduced that although the network's initial response (y_1) is still below the target range, it is significantly increased and hence closer to the target values when compared to figure 6 above. This is a good indication that y , the trained model will thus closely approximate the target values much better and hence will be better trained.

Trained output results

The following discussion gives a brief description of the trained output results against the target ethanol values (in the inputs). One cannot expect to have negative values for temperature, sugar and yeast concentrations, but the reason the x-axis extends to the negative space (Figures D.1 and D.2) is due to that the ethanol concentration is plotted against the coded values of the three variables (that is, temperature, sugar and yeast) as opposed to the real values. Refer to figure D.1 for a plot of the ethanol produced against the real values for the input variables, temperature, and sugar and yeast concentrations.

In all cases, it can be seen that the crosses fall inside the open circles illustrating a good match between the neural network model and the target values for the ethanol concentration. This means that the ANN has been successful in training the data generated and can now be used to forecast the ethanol concentration at varied temperature, sugar and yeast concentrations. This in turn can allow for the optimisation of the ethanol produced in the fermentation process, varying the temperature, sugar and yeast concentration with the aid of a multivariable optimisation method.

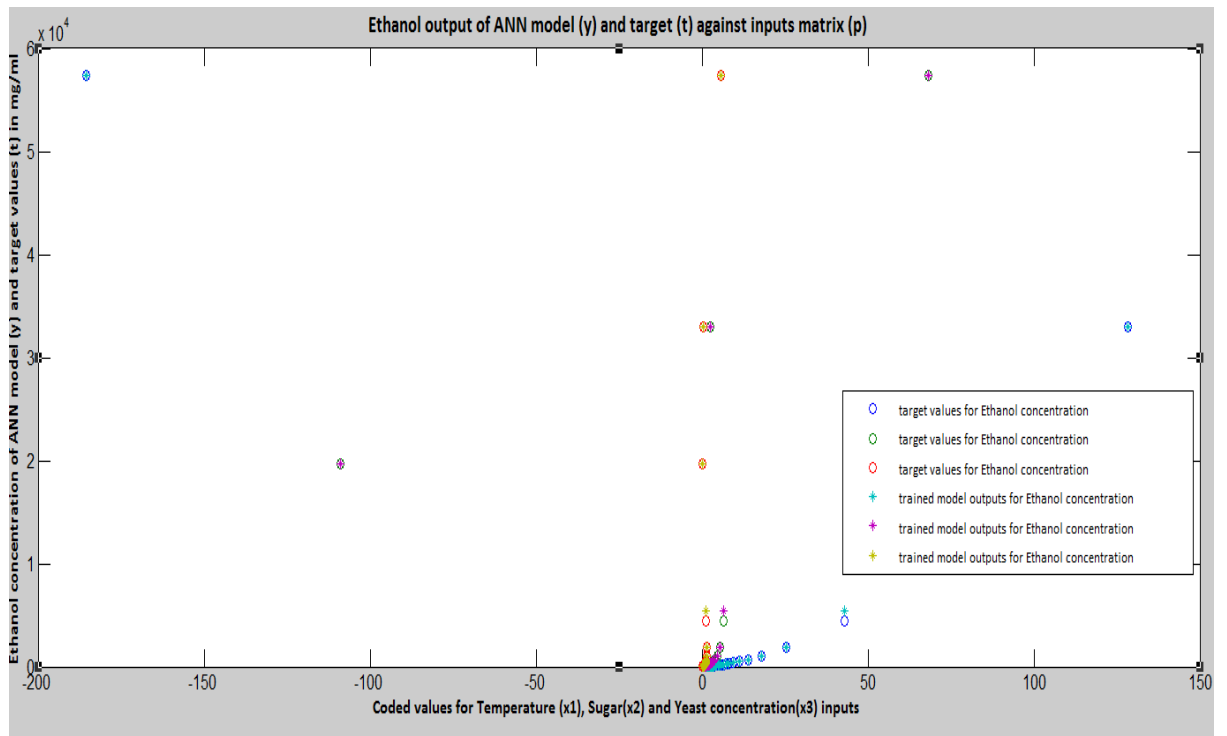


Figure D.3: Trained model output (y) versus target entries (t)

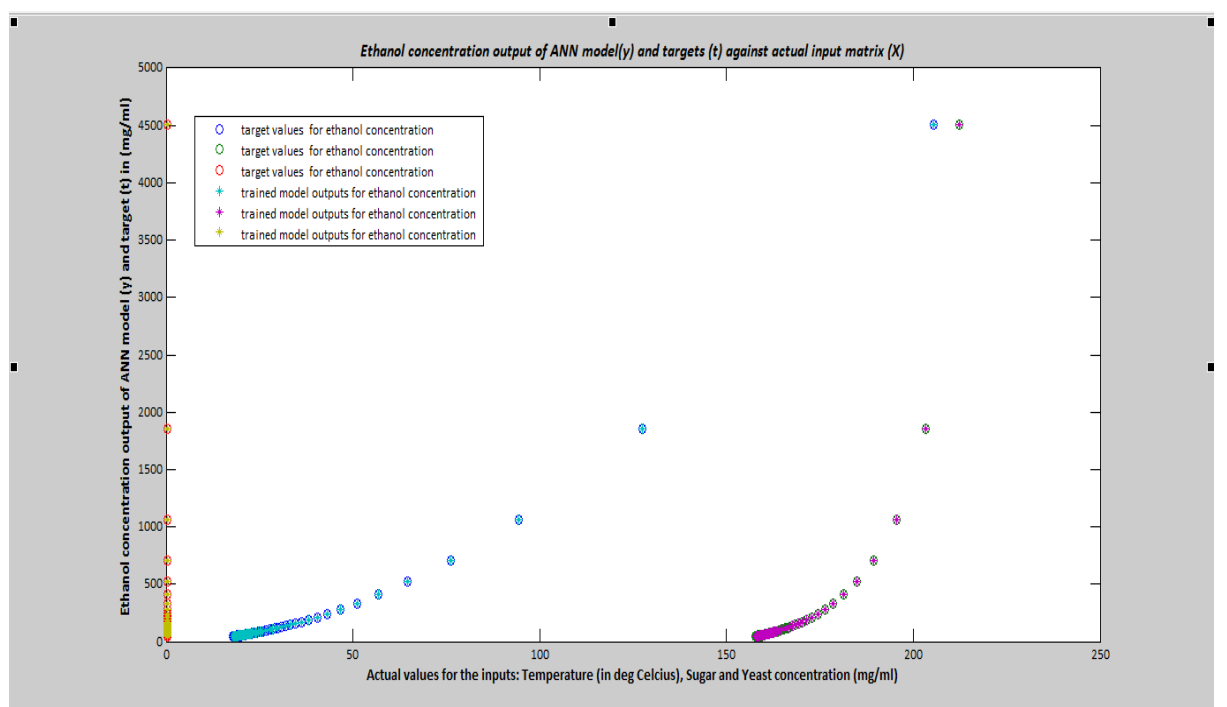


Figure D.4: Trained model output(y) versus target entry (t) against the temperature ($T^{\circ}\text{C}$), sugar and yeast concentration (mg/ml) inputs

The following equation was used to transform the coded values for temperature, sugar and yeast concentrations to the real values (Box and Hunter, 1957):

$$\text{Real Value} = \frac{\text{Upper Case} + \text{Lower Case}}{2} + \frac{\text{Upper Case} - \text{Lower Case}}{2} \times \text{Coded Value} \quad (\text{D.1})$$

Where,

- Subscripts =1, 2 and 3 refer to the temperature, sugar concentration and yeast concentration variables respectively.
- Uppercase is the actual(real) value for one of the three above mentioned process variables
- Lowercase is the coded value for one of the variables mentioned above
- Uppercase represent the actual high (maximum) and low (minimum) values of the variables and
- Lowercase represent the coded high (maximum) and low (minimum) values for the variables

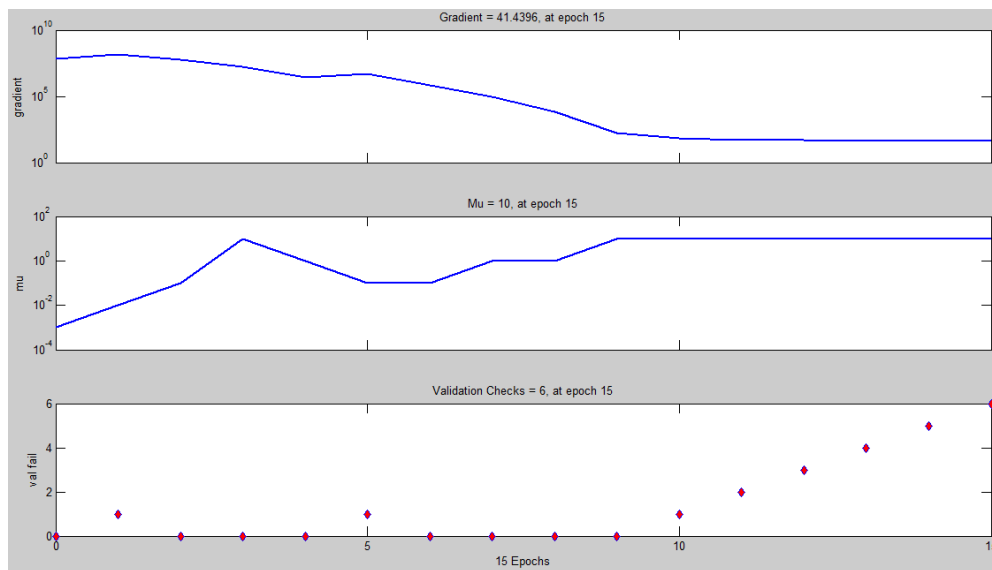


Figure D.5: ANN performance graphs

MATLAB Code for training, validating and testing

```
clc
% Model for generating data for Neural network developed from
experiment.
```

```

u=[5.7:0.1:8]% range of values of u that will give values of x1,x2
and x3 within the limits of -1.682 to 1.682
x1= (-6.00.*(1.70-u).*(0.81-u))./(2.57-6.07.*u+4.38.*(u.^2)-
(u.^3))%calculates value of x1
x2=(-2.45.*(1.89-u).*(0.81-u))./(2.57-6.07.*u+4.38.*(u.^2)-
(u.^3))%calculates value of x2
x3=(-2.00.*(1.89-u).*(1.70-u))./(2.57-6.07.*u+4.38.*(u.^2)-
(u.^3))% calculates value of x3
y=31.16+11.99.*x1+4.90.*x2+3.99.*x3+1.88.*(x1.^2)+1.70.*(x2.^2)+0.
81.*(x3.^2)+0.46.*x1.*x2+5.66.*x1.*x3-4.68.*x2.*x3% y is the
response
X1H=22;X1L=7;x1H=1.682;x1L=-1.682;
X2H=170; X2L=140;x2H=1.682;x2L=-1.682;
X3H=0.5;X3L=0.1;x3H=1.682;x3L=-1.682;
X1=X1H-(((X1H-X1L).*(x1H-x1))./(x1H-x1L))
X2=X2H-(((X2H-X2L).*(x2H-x2))./(x2H-x2L))
X3=X3H-(((X3H-X3L).*(x3H-x3))./(x3H-x3L))
X=[X1;X2;X3]
p=[x1;x2;x3]
target=[y(1) y(2) y(3) y(4) y(5) y(6) y(7) y(8) y(9) y(10) y(11)
y(12) y(13) y(14) y(15) y(16) y(17) y(18) y(19) y(20) y(21) y(22)
y(23) y(24)]
%plot(p,target,'o');
%training the neural network,
%net-Neural Network
%newff- creates a two layer feed forwad network
%sim-presents inputs into a perceptron and to change the weight
%and biases according to error.
%net.trainParam.epochs-states the number of iterations for
training
%net.trainParam.goal-performance goal
net = newff(p,target,10);% creating a two layer feed-forward
network,one hidden layer with ten neurons,p-inputs,t-targets
[trainP,valP,testP]=dividerand(p)%devide input data ,60% of the
inputs are assigned to training set,20% to validation and 20% to
testing
y1 = sim(net,p);%before training simulate network to see if
initial response is close to target.
%plot(p,target,'o',p,y1,'x');% plot of the response before
training
net.trainParam.epochs = 100;%network is trained for up to 100
epochs to an error goal of 0.01
net.trainParam.goal = 0.01;
net = train(net,p,target);%returns a trained neural network
output= (sim(net,p))'
plot(p,target,'o',p,output,'*')%plot of the response after
training
testinput=[1.0025;0.9998;1.0012]%input values which will give the
optimal value of ethanol concentration obtain from using GA
output2=sim(net,testinput)% gives the value of optimal ethanol
concentration using the trained neural network

```

Genetic Algorithm MATLAB script m-file for the function galgorithm

```
function [y]=galgorithm(u)
```

```

x1= (-6.00.*(1.70-u).*(0.81-u))./(2.57-6.07.*u+4.38.*(u.^2)-
(u.^3));
x2=(-2.45.*(1.89-u).*(0.81-u))./(2.57-6.07.*u+4.38.*(u.^2)-
(u.^3));
x3=(-2.00.*(1.89-u).*(1.70-u))./(2.57-6.07.*u+4.38.*(u.^2)-
(u.^3));
y=31.16+11.99.*x1+4.90.*x2+3.99.*x3+1.88.*(x1.^2)+1.70.*(x2.^2)+0.
81.*(x3.^2)+0.46.*x1.*x2+5.66.*x1.*x3-4.68.*x2.*x3;
end

```

m-file for executing the genetic algorithm for the plot to confirm convergence of population (MATLAB 2010,Help files)

```

FitnessFunction = @galgorithm;
numberOfVariables = 1;
[x,Fval,exitFlag,Output] = ga(FitnessFunction,numberOfVariables);
options = gaoptimset('PlotFcns',{@gaplotbestf,@gaplotstopping});
[x,fval] =
ga(FitnessFunction,numberOfVariables,[],[],[],[],[],[],[],options)

```

After running the above GA algorithm file MATLAB returns the following.
Output =

```

    problemtype: 'unconstrained'
      rngstate: [1x1 struct]
    generations: 51
      funccount: 1040
      message: 'Optimisation terminated: average change in the fitness value
less than options.TolFun.'
```

```

The number of generations was : 51
The number of function evaluations was : 1040
The best function value found was : -45.8325
Optimisation terminated: average change in the fitness value less than
options.TolFun.

```

Iteration	Func-count	f(x)	Step-size	First-order optimality
0	2	-45.8352		8.79
1	10	-45.8408	0.000146213	0.0297
2	12	-45.8408	1	0.000468
3	14	-45.8408	1	3.61e-005

Local minimum possible.

fminunc stopped because the size of the current step is less than the default value of the step size tolerance.

<stopping criteria details>

```

x1 =

    1.0025

```

```

x2 =

    0.9998

```

```

x3 =

    1.0012

```

```

fval =

   -45.8408

```

After running the NN script file, MATLAB returns the following (but after several training runs):

```

u =

Columns 1 through 16

    5.7000    5.8000    5.9000    6.0000    6.1000    6.2000    6.3000    6.4000
6.5000    6.6000    6.7000    6.8000    6.9000    7.0000    7.1000    7.2000

Columns 17 through 24

    7.3000    7.4000    7.5000    7.6000    7.7000    7.8000    7.9000    8.0000

x1 =

Columns 1 through 16

    1.5666    1.5267    1.4888    1.4528    1.4184    1.3857    1.3544    1.3245
1.2959    1.2685    1.2422    1.2170    1.1928    1.1696    1.1472    1.1257

Columns 17 through 24

    1.1050    1.0850    1.0657    1.0471    1.0292    1.0118    0.9950    0.9788

x2 =

Columns 1 through 16

    0.6093    0.5945    0.5804    0.5670    0.5542    0.5419    0.5302    0.5190
0.5082    0.4979    0.4880    0.4784    0.4693    0.4605    0.4520    0.4438

Columns 17 through 24

    0.4359    0.4283    0.4209    0.4138    0.4069    0.4003    0.3939    0.3876

x3 =

Columns 1 through 16

    0.4069    0.3988    0.3910    0.3835    0.3763    0.3693    0.3627    0.3562
0.3500    0.3440    0.3382    0.3325    0.3271    0.3218    0.3167    0.3118

Columns 17 through 24

    0.3070    0.3024    0.2979    0.2935    0.2893    0.2852    0.2812    0.2773

y =

Columns 1 through 16

    62.8172    61.8346    60.9089    60.0353    59.2096    58.4283    57.6877    56.9850
56.3173    55.6822    55.0773    54.5007    53.9503    53.4246    52.9219    52.4407

Columns 17 through 24

    51.9798    51.5378    51.1138    50.7065    50.3152    49.9388    49.5766    49.2277

target =

Columns 1 through 16

    62.8172    61.8346    60.9089    60.0353    59.2096    58.4283    57.6877    56.9850
56.3173    55.6822    55.0773    54.5007    53.9503    53.4246    52.9219    52.4407

Columns 17 through 24

    51.9798    51.5378    51.1138    50.7065    50.3152    49.9388    49.5766    49.2277

trainP =

    1.4888    1.4528    1.4184    1.3544    1.3245    1.2685    1.1696    1.1472
1.1050    1.0850    1.0471    1.0292    1.0118    0.9788

```

```

    0.5804    0.5670    0.5542    0.5302    0.5190    0.4979    0.4605    0.4520
0.4359    0.4283    0.4138    0.4069    0.4003    0.3876
    0.3910    0.3835    0.3763    0.3627    0.3562    0.3440    0.3218    0.3167
0.3070    0.3024    0.2935    0.2893    0.2852    0.2773

```

```
valP =
```

```

    1.5666    1.5267    1.2959    1.2422    1.1928
    0.6093    0.5945    0.5082    0.4880    0.4693
    0.4069    0.3988    0.3500    0.3382    0.3271

```

```
testP =
```

```

    1.3857    1.2170    1.1257    1.0657    0.9950
    0.5419    0.4784    0.4438    0.4209    0.3939
    0.3693    0.3325    0.3118    0.2979    0.2812

```

```
output =
```

```

62.8154
62.0539
61.1874
60.2628
58.7457
58.4884
57.7088
56.5934
56.3172
55.6663
55.0474
54.4664
53.4107
53.3772
52.8748
52.3997
51.9421
51.5028
51.0856
50.3850
50.2952
49.9193
49.5666
49.2466

```

```
testinput =
```

```

    1.0025
    0.9998
    1.0012

```

```
output2 =
```

```
63.4884
```

MATLAB Script for calculating the optimal value of Temperature(X1), Sugar Concentration (X2) and yeast Concentation(X3).

```

x1=1.0025;;x2=0.9998;;x3=1.0012;
X1H=22;X1L=7;x1H=1.682;x1L=-1.682;
X2H=170; X2L=140;x2H=1.682;x2L=-1.682;
X3H=0.5;X3L=0.1;x3H=1.682;x3L=-1.682;
X1=X1H-(((X1H-X1L).*(x1H-x1))./(x1H-x1L))
X2=X2H-(((X2H-X2L).*(x2H-x2))./(x2H-x2L))
X3=X3H-(((X3H-X3L).*(x3H-x3))./(x3H-x3L))

```

MATLAB returns

x1 =
18.9701

x2 =
163.9162

x3 =
0.4190