

Supply Chain Complexity and Robustness Analysis

by

Yongyut Meepetchdee

January 2008

A thesis submitted for the degree of
Doctor of Philosophy of the University of London
and for the
Diploma of the Imperial College

Centre for Process Systems Engineering
Department of Chemical Engineering and Chemical Technology
Imperial College of Science, Technology, and Medicine
London, SW7 2BY, United Kingdom

Abstract

This thesis proposes methodologies and analysis frameworks for supply chain efficiency, complexity, and robustness as well as exploring sensitivity, trade-offs, and relationships among them. This thesis investigates the definition and implementation of methods to characterise and analyse complexity and robustness in various aspects of supply chains. We argue that supply chain managers should pay attention to these supply chain performance measures since highly complex supply chains may require more resources and effort to synchronise and coordinate supply chain activities and the robustness metric indicates the long-term serviceability of the supply chain.

Strategic, tactical, and operational aspects of these supply chain performance measures in such areas as logistical network design, vehicle routing, distribution planning, inventory management, and transactional operations are assessed through optimisation and simulation techniques. A Mixed Integer Linear Programming (MILP) model is formulated to design logistical networks with varying robustness requirements at minimum cost, whilst the logistical network complexity is defined as the relative degree of connectivity of the network compared to the minimum spanning tree. Multi-agent based simulation is adopted to simulate distributed supply chain systems and evaluate their performances, including entropic or information-theoretic complexity measure. Quantitative forecasting techniques are employed for demand forecasting. We also characterise and investigate logistical networks and vehicle routes under graph theoretic formalisms to observe average path lengths, clustering coefficients, and degree distributions.

Case studies are performed based on both hypothetical cases and cases from real settings in chemical process industry which allow us to observe several useful relationships between supply chain performance measures. However, the results of this thesis are generic and can be applied to any industry's supply chain.

Acknowledgments

I would like to thank my supervisor Prof. Nilay Shah for giving me the opportunity to conduct research at Imperial College, and for his continuous support, funding, suggestions, invaluable encouragements and excellent supervision throughout my PhD research.

I appreciate Dr. Lazaros Papageorgiou and Dr. Patrick Linke, my PhD thesis examiners, for examining my research work and suggesting improvement on my thesis.

I express my gratitude to Ms. Lynne Smetham and Ms. Christine Bond of INEOS Fluor Ltd. for providing me an opportunity to conduct a challenging case study from the real industrial settings.

I would like to acknowledge my friends and colleagues at CPSE, especially, Kwok Yuen Cheung, Rui Sousa, Yoong Chiang See Toh, Hlynur Stefansson, and Edric Margono for insightful discussions and suggestions about supply chain optimisation and simulation.

My research at CPSE would not have been as smooth and enjoyable without administrative support from CPSE general office staffs, Ms. Senait Selassie and Ms. Cristina Romano.

My appreciation also flows to my Thai friends at Imperial College and in the UK for entertaining me and making me always feel at home.

I am also greatly indebted to the Overseas Research Students Awards Scheme (ORSAS) for their generous financial support during my PhD study.

Finally, I would like to thank my parents, brother, and sisters for their unceasing love and continued support on whatever challenges I undertake.

To my parents,

Mr. Yot and Mrs. Nit Meepetchdee

Table of Contents

Abstract	2
Acknowledgements	3
Table of Contents	5
List of Figures	9
List of Tables	12
1 Introduction	14
1.1 Outline of This Thesis.....	17
2 Literature Review.....	19
2.1 The Supply Chain Overview.....	19
2.2 Supply Chain Management	20
2.2.1 The Bullwhip Effect	21
2.2.2 Supply Chain Partnerships.....	22
2.2.2.1 Vendor Managed Inventory (VMI)	23
2.2.2.2 Supply Contracts	23
2.2.2.3 Distributor Integration (DI)	24
2.2.2.4 Third-Party Logistics (3PL)	25
2.2.3 Design for Logistics.....	25
2.2.4 Information Technology in Supply Chain Management	27
2.2.4.1 Decision Support Systems (DSS)	27
2.3 Supply Chain Modelling	28
2.3.1 Mathematical Programming or Optimisation	28

2.3.2 Simulation.....	30
2.3.3 Types of Supply Chain Problems	31
2.3.3.1 Supply Chain Operation and Execution	31
2.3.3.2 Supply Chain Network and Design	32
2.4 Demand Forecasting.....	34
2.4.1 Characteristics of Forecasts	35
2.4.2 Forecasting Methods.....	36
2.4.2.1 Quantitative forecasting methods	36
2.4.2.2 Qualitative forecasting methods	39
2.4.3 Evaluating Forecasts	41
2.5 Supply Chain Complexity and Complexity Measurement.....	43
2.5.1 The Information-theoretic Measure of Complexity	43
2.5.2 The Chaos-theoretic Measure of Complexity.....	45
2.5.3 The Fitness Measure of Complexity.....	46
2.5.4 Complexity as degrees of connectivity.....	47
2.6 Supply Chain Robustness	48
2.7 Literature Review Summary and Discussion.....	50
 3 Logistical Network Design.....	 52
3.1 Logistical Network Design Concept.....	52
3.2 Mathematical Model.....	57
3.3 A Case Study and Results	62
3.4 Impact of Transportation Economy of Scale	66
3.5 Result of a model with economy of scale in transportation	68
3.6 Network topological properties	69
3.7 Managerial Implications	72
3.8 Summary and Discussion.....	73
 4 Vehicle Routing Networks and Complexity	 74
4.1 Introduction	74
4.2 Problem description	76
4.3 Vehicle routing procedure	77
4.4 Complexity measure	78
4.5 Degree distribution.....	79

4.6	Computational experiments: A case study	80
4.7	Simulation results	81
4.7.1	Results for First Demand Profile	82
4.7.2	Results for Second Demand Profile.....	86
4.8	Degree Distribution in Vehicle Routing Network	91
4.8.1	Probability Distribution Plots	92
4.8.1.1	Degree Distribution Plots for First Demand Profile.....	93
4.8.1.2	Degree Distribution Plots for Second Demand Profile.....	96
4.9	Statistical measures of degree of vertex.....	99
4.10	Managerial Implications	102
4.11	Summary and Discussion	103
5	Multi-Agent Based Supply Chain Modelling and Operational Complexity	104
5.1	Introduction	104
5.2	Modelling framework and agents for supply chain system	106
5.3	Supply chain system performance measures.....	108
5.3.1	Supply Chain Efficiency.....	108
5.3.2	Supply Chain Complexity	108
5.3.3	Fill Rate	110
5.4	Implementation.....	110
5.4.1	Simulation Platform.....	110
5.4.2	Supply Chain Model.....	111
5.4.3	Supply Chain Management Policy	111
5.5	A Case study and Results.....	115
5.5.1	Retailers' Performance	116
5.5.2	Warehouse Performance.....	120
5.6	Summary and Discussion	124
6	Demand Forecasting.....	125
6.1	A case study with INEOS Fluor Ltd.	125
6.1.1	Exponential smoothing technique	126
6.1.1.1	Simple Exponential Smoothing.....	126

	6.1.1.2 Linear Exponential Smoothing (Holt's) Method.....	127
	6.1.1.3 Linear and Seasonal Exponential Smoothing (Winters's) Method	128
	6.1.2 Time Series Decomposition technique.....	128
	6.1.2.1 Decomposition Series with Seasonality only	129
	6.1.2.2 Decomposition of Series with Seasonality and Trend.....	129
	6.1.2.3 Decomposition of Series with Seasonality, Trend, and Cycle	131
	6.2 Forecasting Results	132
	6.2.1 Comparison with INEOS Fluor's Original Forecasts.....	137
	6.3 Residual Analysis	140
	6.4 Summary and Discussion	143
7	Distribution Planning	144
	7.1 INEOS Fluor Ltd. Distribution Infrastructure.....	145
	7.2 A Simulation Model.....	146
	7.3 A case study.....	147
	7.4 Costs, Service Levels, and Complexity Calculation.....	151
	7.5 Scenarios Setting.....	153
	7.6 The simulation results	157
	7.7 Summary and Discussion.....	163
8	Conclusions and Further Work	165
	8.1 Conclusions	166
	8.2 Recommendations for Future Work.....	168
	References	172
	Appendix A: Distance matrix for a case study in chapter 3	184
	Appendix B: Derivation of the Equations for the Slope and Intercept for Regression Analysis.....	185

List of Figures

2.1	Typical Supply Chain Structure	20
2.2	Increasing Variability of Orders up the Supply Chain	21
3.1	Customer zone geography	62
3.2	Customer demand distribution.....	62
3.3	Logistical network structure (no robustness requirement)	63
3.4	Logistical network structures for various robustness requirements.....	64
3.5	Relationships of Robustness versus Efficiency and Complexity.....	65
3.6	Different logistical network structures with topological or redundancy robustness	66
3.7	Unit transportation cost rate vs. Shipping Quantity.....	67
3.8	Logistical network of model with transportation economy of scale.....	68
3.9	Degree distributions of logistical networks	70
4.1	Customer zone geography	76
4.2	Average monthly demand profile for each customer zones.....	80
4.3	Example of vehicle routing (centralised distribution)	82
4.4	Average complexity vs. Logistics configuration (1st demand profile).....	83
4.5	Complexity level vs. Planning horizon (1st demand profile)	84
4.6	Number of trucks required for each day in configuration A with small truck	85
4.7	Average complexity vs. Logistics configuration (2nd demand profile)	87
4.8	Complexity level vs. Planning horizon (2nd demand profile).....	88
4.9	Complexity Level vs. Number of Trucks in fixed route vehicle routing.....	90
4.10	Illustration of edges direct toward and direct away from nodes.....	91

4.11	Degree distribution plots of degree of vertex for large truck (1st demand profile)	93
4.12	Degree distribution plots of degree of vertex for medium truck (1st demand profile)	94
4.13	Degree distribution plots of degree of vertex for small truck (1st demand profile)	95
4.14	Degree distribution plots of degree of vertex for large truck (2nd demand profile)	96
4.15	Degree distribution plots of degree of vertex for medium truck (2nd demand profile)	97
4.16	Degree distribution plots of degree of vertex for small truck (2nd demand profile)	98
4.17	Level of consistency with the power-law degree distribution	99
4.18	Mean and Variance of degree of vertex (1st demand profile)	100
4.19	Mean and Variance of degree of vertex (2nd demand profile)	101
5.1	Supply Chain Network Structure.....	111
5.2	Market agent and Retailer agent's supply chain policy	112
5.3	Example of Warehouse Inventory Projection when $I_0 = B_T$	113
5.4	Warehouse agent's supply chain policy	114
5.5	Manufacturer agent's supply chain policy and raw material agent	115
5.6	Effect of RFOI to retailers' performance.....	117
5.7	Relationships between retailers' supply chain performance measures.....	119
5.8	Effect of WFOI to warehouse performance.....	121
5.9	Relationships between warehouse supply chain performance measures	123
6.1	Forecasting results from Simple Exponential Smoothing	133
6.2	Forecasting results from Linear Exponential Smoothing	133
6.3	Forecasting results from Linear and Seasonal Exponential Smoothing	134
6.4	Forecasting results from TSD (Season only).....	134
6.5	Forecasting results from TSD (Season and Trend).....	135
6.6	Forecasting results from TSD (Season, Trend, and Cycle)	135
6.7	Cyclical pattern from TSD (Season, Trend, and Cycle)	136
6.8	Mean Absolute Percentage Error bar chart.....	137

6.9	Average errors and standard deviation of errors bar chart.....	137
6.10	Product A sales forecasts for market X	139
6.11	Product A sales forecasts for market Y	139
6.12	Histogram of residual errors from Simple Exponential Smoothing	140
6.13	Histogram of residual errors from Linear Exponential Smoothing	141
6.14	Histogram of residual errors from Linear and Seasonal Exp. Smoothing	141
6.15	Histogram of residual errors from TSD (Season only).....	141
6.16	Histogram of residual errors from TSD (Trend and Season).....	142
6.17	Histogram of residual errors from TSD (Trend, Season, and Cycle)	142
7.1	Diagram of product supply for product A and product B	145
7.2	Diagram of the simulation mechanism	147
7.3	User interface on NetLogo	148
7.4	Format of inputs in a text file	151
7.5	Distribution costs in different scenarios	157
7.6	Distribution costs per unit delivered in different scenarios	159
7.7	OTIF rates in different scenarios	160
7.8	Fill rate in different scenarios.....	161
7.9	Complexity levels in different scenarios	162

List of Tables

3.1	Logistical network performance measurement table	63
3.2	Performance measurement of network with scale economy	68
3.3	Summary of topological properties of logistical networks	71
4.1	Logistics Configuration	77
4.2	Summary of complexity level (1st demand profile)	82
4.3	Statistics of number of trucks required and complexity level (1st demand profile)	86
4.4	Summary of complexity level (2nd demand profile).....	87
4.5	Statistics of number of trucks required and complexity level (2nd demand profile)	89
4.6	Complexity level in fixed route vehicle routing	90
4.7	Mean and Variance of degree of vertex (1st demand profile)	100
4.8	Mean and Variance of degree of vertex (2nd demand profile).....	101
5.1	Generalised replenishment quantity for warehouse's replenishment plan....	113
5.2	List of parameters setting for the case study	116
6.1	Summary of forecasting accuracy	136
6.2	Absolute Percentage Error of INEOS's forecasts and TSD forecasts	138
7.1	Parameters input for a case study simulation	154
7.2	Monthly demands for product A and product B	155
7.3	Details of Isotanks in transit at the beginning of the simulation.....	155
7.4	Details of Isotanks that will be available in future date at each location.....	156

7.5 Details of Isotanks that will be lost in future date at each location 156

7.6 Details of five different backhaul strategies 156

7.7 Average cost and Standard Deviation..... 157

7.8 Average cost per unit delivered in different scenarios 158

7.9 Average OTIF rate and Standard Deviation 159

7.10 Average Fill rate and Standard Deviation 160

7.11 Average Complexity level and Standard Deviation 162

Chapter 1

Introduction

In the 1980s companies discovered new manufacturing technologies and strategies that allowed them to reduce costs and compete better in different markets. In the last few years, however, it has become clear that many companies have reduced manufacturing costs as much as is practically possible. Many of these companies are discovering that effective supply chain management is the next step they need to take in order to increase profit and market share.

Supply Chain Management is a set of approaches utilised to efficiently integrate suppliers, manufacturers, warehouses, and stores, so that merchandise is produced and distributed at the right quantities, to the right locations, and at the right time, in order to minimise systemwide costs whilst satisfying service level requirements.

Even though the term “Supply Chain Management” has just been coined quite recently, the concept and the importance of supply chain management have been in existence for more than a hundred years. In a book “A Century in Oil: The Shell Transport and Trading Company 1897-1997”, Stephen Howarth (1997), the author, vividly described as follows the magnificently and secretly designed supply chain strategy that helped Shell, a fledgling company in 1892, toppled Standard Oil, a world market dominant American firm, in terms of market share.

“... First, the source of oil would have to be utterly reliable, and priced to reflect the refiner’s savings on the cost of tinplate and cases. Second, not just one ship but a fleet would have to be built simultaneously, to ensure continuity of supply to the markets. Third, the ships would have to be much bigger than the Nobel’s - bigger, indeed, than the *Glückauf*, and better,

because the most economical route from the Black Sea to the Far East was via the Suez Canal, and there, the regulations governing the transport of oil in bulk were so stringent that no one (not even Standard, which had tried) had ever received permission to take such a cargo through. Fourth, it would be pointless to try the experiment in only one country, or here and there through the region; that would invite piece-meal destruction by Standard. Therefore, a complete network both of storage, distribution and sales would have to be in place before the first shipment of oil was sent ...”

Very similar strategies as above are still, indeed, being employed in today’s business. Over the past few decades, Supply Chain Management has become an important focus of competitive advantage for firms and organisations. Well-managed supply chains will significantly improve cost effectiveness and sustain long term profitability. Managing the supply chain requires a wide range of methods and principles. One way of classifying Supply Chain Management (SCM) analysis is to divide the area into Supply Chain Design (SCD) and Supply Chain Operation or Execution (SCE).

At the core of Supply Chain Design is Strategic Network Design which allows planners to pick the optimal number, location, and size of warehouses and/or plants; to determine optimal sourcing strategies, that is, which plant/vendor should produce which product; and to determine the best distribution channels, that is, which warehouses should service which customers. The objective is to minimise total costs, including sourcing, production, transportation, warehousing, and inventory, by identifying the optimal trade-offs between the number of facilities and service levels. The planning horizon for these systems is typically a few months to a few years using aggregated data and long-term forecasts. The method most often used is optimisation. Strategic Network Design is attractive in that its return on investment is usually higher than operational planning and execution. The network design decisions also impact large capital investments and major distribution decisions.

Supply Chain Operation or Execution involves planning and execution of supply chain operations for short-term or medium-term, for example, demand forecasting, production scheduling, inventory management, distribution planning, and vehicle routing, to name a few. These supply chain activities are not only important for manufacturing industry or industries that produce physical goods, but also for service industries such as the airline business or the financial services industry. Supply Chain Operation offers potential strategic advantages to firms that can efficiently manage these activities, especially in terms of cost, service level, and responsiveness.

Supply chain management is ever more important in today's world of business where competition is extremely intense with ever higher customer expectations for products and services. In addition, with globalisation, deregulation, and new opportunities in emerging markets, supply chain networks and operations are becoming more distributed which adds to complexity in the supply chain.

In the past, a major focus among supply chain practitioners has been on efficiency. However, we now realise that when only short-term profitability is taken into consideration, long-term serviceability might be sacrificed. Hence, recently there has been growing attention on *robustness* of the supply chain which is the measure of supply chain serviceability at the present of disruption or uncertainty.

Due to the fact that today's supply chains are becoming increasingly more complex and difficult to manage, this thesis aims at the definition and implementation of methods to characterise and analyse complexity in supply chain design and operations, as well as the robustness of supply chain. It brings together the experience in supply chain modelling with the recently emerging field of robustness and complexity analysis. The main goals of this thesis are:

- Investigating application of robustness and complexity analysis to supply chains
- Identifying methods to quantify robustness and complexity
- Proposing methods to formulate and optimise complex supply chain networks with different levels of efficiency and robustness
- Relating robustness and complexity to other supply chain performance measures
- Investigating whether there is a potential benefit of complexity to improve supply chain performance through scientific management of complexity
- Investigating supply chain improvement opportunity using operations research and management science tools

We will use a number of available mathematical and computational tools such as optimisation, simulation, and time series analysis to design, analyse, and evaluate supply chain systems. Furthermore, we shall apply sensitivity analysis to provide insights into how constraints or different factors affect supply chain performance, and identify major parameters that could significantly improve the performance. The result of the research project will be generic so that it can be applied to any industry's supply chain.

1.1 Outline of This Thesis

This thesis is divided into eight chapters. In this first chapter, we introduce the concepts of supply chain management and its importance in today's business. We illustrate some historical account in supply chain management, classification, and current concerns in supply chain practice. We demonstrate the new challenge of supply chain complexity and robustness that supply chain managers encounter and describe the aims and scope of this thesis to address these challenges.

Chapter 2 provides a review of past and current research relating to supply chain management, design, and analysis. This chapter covers a literature review of supply chain management, supply chain modelling, measures of complexity and robustness, and forecasting techniques. This chapter also presents a basic description of solution techniques for supply chain problems such as optimisation and simulation.

In chapter 3, we propose and develop a logistical network design framework with robustness and complexity considerations. A mathematical model formulated as a Mixed Integer Linear Programming (MILP) is constructed based on the conceptual framework and applied to a hypothetical case study with varying robustness requirements. The network complexity is then measured. Furthermore, we introduce a graph-theoretic view of the logistical network and present its topological properties such as average path length, clustering coefficient, and degree distribution.

Chapter 4 builds upon the logistical networks designed in chapter 3 to present an approach to evaluate complexity in vehicle routing operations. We consider vehicle routing as a system where the states of the system can be defined as different routes of a vehicle. The system complexity is calculated based on an entropic complexity approach. We also investigate vehicle routing network topological properties and compare the properties of alternative network designs.

As the discrete and distributed nature of the supply chain makes it more difficult to evaluate supply chain operation analytically, in chapter 5, we present the assessment of supply chain performance through a simulation method based on a multi-agent based modelling concept. We develop a modelling framework, define supply chain policies, and create a simulation model. We perform simulation on our hypothetical case study to evaluate supply chain performance in

supply chain environments under different conditions of operational complexity. We also relate the operational complexity to other supply chain performance measures.

Chapter 6 discusses the case study at INEOS Fluor Ltd., a world leading fluorine-based products manufacturing company, which experiences highly fluctuating monthly product demand patterns throughout the year. Therefore, we introduced the INEOS Fluor management to quantitative forecasting techniques. In this chapter, we develop quantitative forecasting models based on INEOS Fluor's historical sales data and compare forecasting accuracy to current INEOS Fluor's current forecasting method. We also perform residual analysis to gain insights from distribution of residual errors.

Chapter 7 addresses another case at INEOS Fluor where the management is contemplating a distribution strategy by identifying appropriate levels of product's container backhauling. As an aid to a supply chain manager, in this chapter, we develop an agent-based simulation tool to reveal distribution costs, service levels, and complexity of different distribution management alternatives. Forecasted demands such as those in chapter 6 will be used as one of the inputs for the simulation model.

The final chapter, chapter 8, concludes the findings from our work in this thesis. We summarise the overall results and suggest possible directions for further investigation of supply chain complexity and robustness analysis.

Chapter 2

Literature Review

In this chapter, a general background of the study on supply chain management, supply chain modelling, supply chain networks, forecasting, and robustness and complexity measurement is presented. This is followed by a review of relevant literature in the above areas. In general, the primary objective of supply chain management is either to maximise profit or to minimise cost whilst still satisfying a required set of service levels. Therefore, in this thesis, we use profit maximisation or cost minimisation interchangeably in the context of the problems to indicate optimal supply chain performance.

2.1 The Supply Chain Overview

A supply chain consists of all parties involved, directly or indirectly, in fulfilling a customer request. The supply chain not only includes the manufacturer and suppliers, but also transporters, warehouses, retailers, and customers themselves. Within each organisation, such as a manufacturer, the supply chain includes all functions involved in receiving and filling a customer request. These functions include, but are not limited to, new product development, marketing, operations, distribution, finance, and customer service (Chopra, 2004).

The typical supply chain structure in Figure 2.1 below illustrates how components in the supply chain connect to each other.

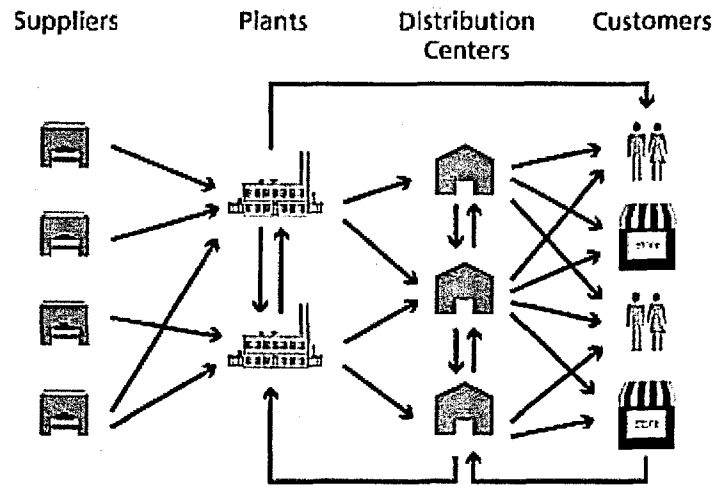


Figure 2.1: Typical Supply Chain Structure

Arrows in the figure demonstrate flows of materials and goods in the supply chain. However, the supply chain not only involves in flow of physical products but also flow of information and communications. In a typical supply chain, raw materials are procured at one or more factories, shipped to warehouses or distribution centres for intermediate storage, and then shipped to retailers or customers. The supply chain, which is also referred to as the logistics network, consists of suppliers, manufacturing centres, warehouses, distribution centres, and retail outlets, as well as raw materials, work-in-process inventory, and finished products that flow between the facilities (Simchi-Levi *et al.*, 2003).

2.2 Supply Chain Management

This section presents the literature sources from a management perspective, which qualitatively and/or quantitatively describes concepts and examples of supply chain management.

In order to create and sustain competitive advantages for a supply chain, Muckstadt *et al.* (2001) proposed that operational uncertainty must be reduced and dealt with explicitly by all supply chain partners. To cope with uncertainty, they proposed guidelines for supply chain design with five principles of supply chain management excellence for the effective design and execution of supply chain systems that must be followed in concert. The five principles are: know the customer, construct a lean supply chain organisation, build tightly coupled

information infrastructures, build tightly coupled business processes, and construct tightly coupled decision support systems.

In addition, frequently, opportunity for supply chain improvement lies at the interface between supply chain partners. In his world renowned book, *Competitive Strategy*, Porter (1980) developed strategies toward both buyers and suppliers. Selecting the right buyer will enhance the profitability of a company. The criteria for buyer selection are buyer's purchasing needs versus company capabilities, buyer's potential growth, buyer's bargaining power, price sensitivity of buyer, and cost of servicing buyer. Similarly, important purchasing strategies are avoiding switching costs, promoting standardisation, qualifying alternative sources, to name a few.

2.2.1 The Bullwhip Effect

In recent years many suppliers and retailers have observed that whilst customer demand for specific products does not vary much, inventory and back-orders fluctuate considerably across their supply chains. For instance, in examining the demand for Pampers disposal diapers, executives at Procter & Gamble noticed that even though retail sales of the product were fairly uniform, distributors' orders placed on the factory fluctuated significantly. Furthermore, when they tracked the order pattern of the factory to raw materials suppliers, they found even higher fluctuations. This increase in variability as one travels up in the supply chain is referred to as the *Bullwhip Effect*. Figure 2.2 provides a graphical representation of orders, as a function of time, placed by different facilities.

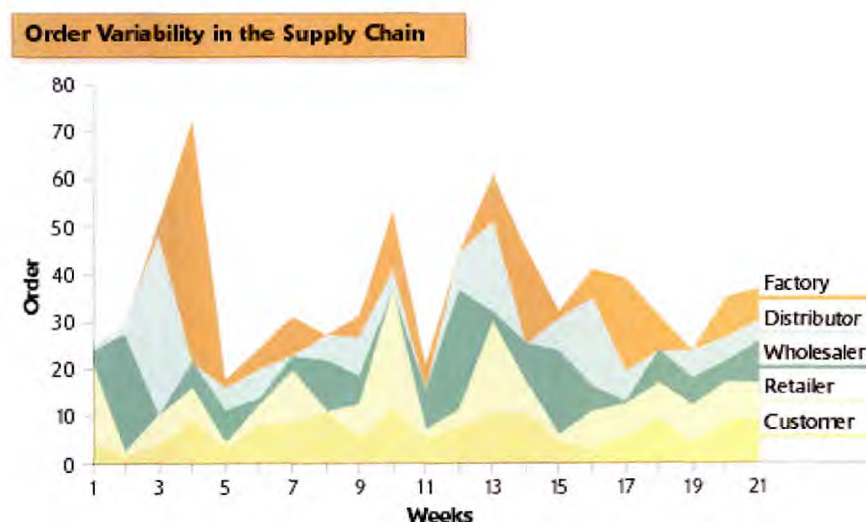


Figure 2.2: Increasing Variability of Orders up the Supply Chain

(Simchi-Levi *et al.*, 2003, pp. 103)

The Bullwhip Effect implies that the manufacturer and supplier who only observe their immediate order data will be misled by the amplified demand patterns, and this has serious cost implications. For example, the manufacturer incurs excess raw materials cost due to unplanned purchases of supplies, additional manufacturing expenses created by excess capacity, inefficient utilisation and overtime, excess warehousing expenses and additional transportation costs due to inefficient scheduling and premium shipping rates.

Lee *et al.* (1997) analysed four sources of the bullwhip effect: demand signal processing, rationing game, order batching, and price variations, and identified strategies to alleviate the detrimental impact of the bullwhip effect. These strategies require among other things: information sharing of sell-through and inventory status data, coordination of orders across retailers and simplification of the pricing/promotional activities of the manufacturer. They claim that demand distortion may arise as a result of local optimising behaviours by players in the supply chain. They also noted that although this shift in attitude is desirable to the *manufacturer*, benefit sharing is normally needed to provide an incentive to the *retailer* to provide the manufacturer with the data.

2.2.2 Supply Chain Partnerships

The main cause of the bullwhip effect described above can be collectively attributed to the lack of coordination or, in other word, local optimising behaviours. When the system is not coordinated – that is, each facility in the supply chain does what is best for that facility – the result is *local optimisation*. Each component of the supply chain optimises its own operation without due respect to the impact of its policy on other components in the supply chain. The overall system could be better off if each component does not take a myopic view, but instead cooperates to optimise the entire system. This approach is termed as *global optimisation*. Although this could lead to an increase in costs in one system, larger decreases will occur elsewhere so that the overall cost is lower. Global optimisation is obvious if the entire system is owned by a common owner. However, even if there is no such a common owner incentives such as benefit sharing can be introduced to provide motives for all components to cooperate and optimise globally.

To achieve global optimality whilst fulfilling local satisfaction, Gjerdrum *et al.* (2001a) developed a mathematical programming formulation so as to achieve solutions of similar quality as global optimisation but with fair profit distributions based on the game theoretical Nash-type objective function. Model decision variables include inter-company transfer prices,

production and inventory levels, resource utilisation, and flows of products between echelons, subject to a deterministic sale profile, minimum profit requirements for each enterprise, and various resource constraints. The application is demonstrated through case studies. The proposed method has proved to be a promising way to share benefit among all players in the supply chain, whilst attaining the highest possible profit for the overall system. The approach assumes that all relevant information is available for the entire overall supply chain optimisation. However, the authors also noted that it is still quite common in practice that players are reluctant to make all information available to their partners.

Further cooperation can also be achieved by engaging one or more types of *Supply Chain Partnerships* as discussed below.

2.2.2.1 Vendor Managed Inventory (VMI)

In vendor managed inventory, the manufacturer manages the inventory of its product at its customer's site (e.g. a retailer outlet), and therefore determines for itself how much inventory to keep on hand and how much to ship to the customer in every period. Therefore, in VMI the manufacturer does not rely on the orders placed by its customer, thus avoiding the bullwhip effect entirely. VMI requires the customer to share demand information with the manufacturer to allow them to make inventory replenishment decisions. VMI can allow a manufacturer to increase their profits by more efficient resource utilisation, as well as to increase profits for the entire supply chain. The customer may request the manufacturer to own the inventory until the goods are sold to prevent the manufacturer from moving too much inventory to the customer. As a result, the customer's cost decreases and profit increases.

Schonberger (1996) discussed three types of collaboration activities; multi-echelon planning, planning with immediate partners, and internal multifunctional planning. Of the three types of collaboration, multi-echelon agreements (e.g. vendor-managed inventory) appear to have the greatest potency. For example, by 1992 Kmart had developed 200 suppliers as its VMI partners. Wal-Mart has also publicly announced that it would cease doing business with distributors and would deal only with manufacturers.

2.2.2.2 Supply Contracts

Optimal supply chain performance requires the execution of a precise set of actions. Unfortunately, those actions are not always in the best interest of all the members in the supply

chain. Arranging Supply Contracts is one of the techniques to motivate each member in the supply chain to align their objectives such that the entire system is optimised. These contracts address issues that arise between a buyer and a supplier. In a supply contract, the buyer and supplier may agree on

- Pricing and volume discounts
- Minimum and maximum purchase quantities
- Delivery lead times
- Product or material quality
- Product return policies

Supply contracts are very powerful tools that can be used for far more than to ensure adequate supply of, and demand for, goods. In fact, a careful design of these contracts can achieve almost the same profit as the profit in global optimisation, hence, a motivation for valuing supply chain contracts.

Kamrad and Saddique (2004) analysed and valued supply contracts in an environment characterised by exchange rate uncertainty, order-quantity flexibility, supplier-switching options, reaction options, and profit sharing. They adopted a real-options approach to the analysis and valuation of such contracts, and viewed the producer's supplier selection set as a portfolio of risky assets employing basic diversification concepts from portfolio theory and provide a unique frame work for risk reduction.

2.2.2.3 Distributor Integration (DI)

Modern information technology leads to a capability in which distributors are integrated so that the expertise and inventory located at one distributor is available to the others. Distributor Integration (DI) can be used to address both inventory-related and service-related issues. In terms of inventory, DI can be used to create a large pool of inventory across the entire distributor network, lowering total inventory costs while raising service levels. Similarly, DI can be used to meet a customer's specialised technical service requests by steering these requests to the distributors best suited to addressing them.

In implementing DI, Narus and Anderson (1986) advised that, to run businesses smoothly, manufacturers should treat distributors like their partners. For manufacturers to effectively plan and implement industrial distributor programs they must: gain a deep

understanding of distributor requirements, build working partnerships with distributors, and actively manage these partnerships.

2.2.2.4 Third-Party Logistics (3PL)

Business practices nowadays are becoming more complex, and frequently a company does not have enough financial and managerial resources to implement all these practices. Even if a firm has the available resources to perform particular task, another firm in the supply chain may sometimes be able to perform that task much more efficiently and cheaply. Third-party logistics is the use of an outside company for all or part of the firm's materials management and product distribution functions. The most frequently cited benefit of using 3PL providers is that it allows a company to focus on its core competencies. With corporate resources becoming increasingly limited, it is often difficult to be an expert in every facet of the business. 3PL also provides technological and other flexibilities in a more cost-effective way. However, if logistics is one of the core competencies of a firm, it makes no sense to outsource these activities to a supplier who may not be as capable as the firm's in-house expertise.

In investigating 3PL success factors, Leahy *et al.* (1995) reported the results from a survey, using a questionnaire, of leading third-party providers in the US in an attempt to address factors for a successful or effective third-party relationship. Customer orientation and dependability are the two factors that received the highest importance ratings; sharing human resources was assigned the lowest importance ratings. Additional analysis also indicated that the respondents generally agree on the factors important to successful third-party relationships.

2.2.3 Design for Logistics

In the 1980s, management realised that product and process design were key product cost drivers, and that taking the manufacturing process into account early in the design process was the only way to make the manufacturing process efficient. Thus, the concept of design for manufacturing (DFM) was born. Recently, a similar transformation has begun in the area of supply chain management. In the last few years, managers have started to realise that by taking supply chain concerns into account in the product and process design phase, it becomes possible to operate a much more efficient supply chain. The concept is known as Design for Logistics (DFL) or Design for Supply Chain Management (DFSCM), whose key techniques and considerations are

- Economic packaging and transportation
- Concurrent and parallel processing
- Standardisation and postponement

Each of these components addresses the issue of inventory or transportation costs and service levels in complementary ways.

Swaminathan (2001) asserted that product variety is increasingly important in many businesses. Along with product variety comes increased uncertainty and a number of operational challenges. These challenges can be met by using various types of standardisation – product, part, procurement, and process. The ability for a firm to use these standardisation approaches depends in part on the degree to which it can modularise its products and processes, in part on what it is trying to achieve for its customer and in part on the costs associated with standardisation. He also presented a framework that assists in choosing among these options.

Feitzinger and Lee (1997) highlighted the success of applying standardisation and postponement concepts at Hewlett-Packard to manufacture and deliver customised products quickly and at a low cost under highly uncertain demand. The key to mass-customising effectively is postponing the task of differentiating a product for a specific customer until the latest possible point in the supply network. In doing so, a company must achieve modular product design, modular process design, and agile supply networks. Standardisation may come with additional materials costs. This has to be compared with reductions in inventory cost and possible increases in sales.

Yang and Burns (2003) discussed the implications of postponement for the supply chain with a view to extending the significance of postponement towards the perspective of a holistic supply chain context. They highlighted implications of postponement in relation to the decoupling point, supply chain integration, control of the supply chain and capacity planning issues. They also noted that regardless of industry, using postponement as a supply chain strategy (e.g. colouring paints at the filling stage) is not easy. Companies must understand the costs and benefits of postponement, as well as understand how to tackle obstacles in the way to an effective adoption.

Postponing an action is one of the strategies to handle uncertainty in risk management in general, not only in supply chain management. Bernstein (1996) remarkably explained the value of postponement using an example from Hamlet as follows;

"... Hamlet complained that too much hesitation in face of uncertain outcomes is bad because *'the native hue of resolution is sicklied o'er with the pale cast of thought... and enterprises of*

great pith and moment... lose the name of action. Yet once we act, we forfeit the option of waiting until new information comes along. As a result, not-acting has value. The more uncertain the outcome, the greater may be the value of procrastination. Hamlet had it wrong: he who hesitates is halfway home..."

Lee and Sasser (1995) discussed their experience in evaluating two design alternatives for power supply in Hewlett-Packard's Rainbow product, a computer peripheral device. An inventory and cost model was used in the evaluation. The objective was to quantify the cost difference on a per unit basis between universal (or standardised) and dedicated power supply. It was found that, even under conservative assumptions, the cost benefits from universal power supply are significantly greater than the increase in material and manufacturing cost incurred.

van Hoek *et al.* (1999) studied the implementation of postponement strategies in European supply chains. They provided a table of factors favouring postponement such as feasibility to decouple primary and postponed operation, high commonality of modules etc. versus impact of postponement. A case study was conducted on implementation of postponement strategy through interviews and direct observation in four companies in different industries. The results revealed that the majority of factors favouring postponement match the operational characteristics of the four companies, and the companies can either reduce operating costs or improve lead times, or both by implementing a postponement policy.

2.2.4 Information Technology in Supply Chain Management

Information is crucial to the performance of a supply chain because it provides the basis upon which supply chain managers make decisions. Thus, information technology plays a very critical role in supply chain management. It is an enabler for implementation of many supply chain management strategies. For example, supply chain integration, vendor managed inventory (VMI), distributor integration (DI), third-party logistics (3PL) etc., are all hardly effectively executed without the ability to collect and process data quickly and cheaply throughout the supply chain.

2.2.4.1 Decision Support Systems (DSS)

A Decision Support System (DSS) is a set of computerised tools that assist and provide relevant information to a decision maker for making a decision under particular circumstances. A DSS could range from spreadsheets, in which users perform their own analysis, to expert systems, which attempt to incorporate the knowledge of experts in various fields and suggest possible

alternatives. DSS analysis tools and techniques can be as simple as data queries to sophisticated mathematical modelling, depending on types of problems and needs of users. Many DSS also provide user-friendly reports and graphical presentation. Some types of supply chain decision that can utilise DSS are:

- Logistics Network Design
- Supply Chain Master Planning
- Operational Planning
- Inventory Management
- Transportation Planning
- Production Scheduling
- Material Requirements Planning
- Operational Execution

Given the decision support tools available, Bender (2000) discussed misconceptions of adopting decision support systems in supply chain management. He argued that an organisation does not need to have an enterprise resource planning (ERP) system in place before it implements decision support systems. He also pointed out that a rule-based decision support system is not the best way to achieve optimal supply chain performance. A preferable approach is a mathematical one that can efficiently represent all the options before a decision maker and then identify the best possible one.

2.3 Supply Chain Modelling and Analysis

Formulating supply chain management strategy and operational procedures is usually neither intuitive nor easily obtained by simple enumeration. Often, it is necessary to describe or model characteristics, constraints, and objectives of the supply chain in mathematical terms and then find the best possible solution or average performance measurement. The most widely used techniques in supply chain modelling are:

2.3.1 Mathematical Programming or Optimisation

This technique is used to find potential optima or the best possible solution. For example, it may generate the best set of locations for new warehouses, the selection of the best manufacturing process, an efficient route for a truck to take, or an effective inventory policy for a retail store. The optimisation model is described by a set of constraints (equations or inequalities) and one or more objective functions. Major types of optimisation model are:

Linear Programming

Linear Programming is an optimisation model where the mathematical functions appearing in both the objective function and the constraints are all linear functions. In addition, all decision variables are allowed to have non-integer values.

Integer Programming

In many practical problems, the decision variables make sense only if they have integer values. For example, it is often necessary to assign people, machines, and vehicles to activities in integer quantities. If requiring integer values is the only way in which a problem deviates from a linear programming formulation, then it is an *Integer Programming* problem. If only some of the variables are required to have integer values (the rest can be non-integer), then the model is referred to as *Mixed Integer Linear Programming* (MILP).

Nonlinear Programming

Certain problems may have some degree of nonlinearity. When the objective function or at least one of the constraints is nonlinear, the model will be a Nonlinear Programming model.

Furthermore, if some of the variables are required to have integer values, then the model will be referred to as *Mixed Integer Nonlinear Programming* (MINLP).

A limitation of some optimisation techniques is that they deal with static models and they do not take into account changes over time. The solution techniques for optimisation model can generally be classified into two major sets of algorithms:

- Exact algorithms, which guarantee to find optimal solutions if at least one feasible solution exists. However, these kinds of algorithms may take a long time to run, especially if a problem is very complex. Examples of exact algorithms are Simplex Method, Interior-Point Algorithm (both for Linear Programming), Branch-and-Bound Algorithm (for Integer Programming and MILP), Newton's Method, and Gradient Search Procedure (both for Nonlinear Programming). However, algorithms for nonlinear programming only guarantee local optimum but not global optimum.
- Heuristic algorithms. Heuristics aim to provide sufficiently good, but not necessarily optimal, solutions. They typically run much faster than exact algorithms. A good heuristic will rapidly give a solution that is very close to the optimal solution. Heuristic design often involves a trade-off between the quality of a solution and computational speed. Normally, Linear Programming problems can be solved quickly using exact algorithms. Therefore, heuristic algorithms are usually used in Integer Programming,

MILP, Nonlinear Programming, or MINLP, where they may be used in conjunction with some exact algorithms. Examples of heuristic algorithms are Tabu Search, Simulated Annealing, and Genetic Algorithms.

2.3.2 Simulation

In contrast to optimisation techniques, simulation-based tools can take into account the detailed dynamics of the system and characterise system performance for a given design. Simulation is frequently an effective tool to help analyse problems with some random, or stochastic, elements. In simulation, a model of the process is created on a computer. Each of the random elements of the model is specified with a probability distribution. When the model is run, the computer simulates carrying out the process. Each time a random event occurs, the computer uses the specified probability distribution to randomly decide what happens. Since it is a random model, each time the model is run, the results may be different. Statistical techniques are used to determine the model's average outcome and the variability of this outcome. Also, by varying input parameters, different models and decisions can be compared. Simulation is often a useful tool to understand very complex systems that are difficult to analyse analytically. In particular, a multi-agent based simulation has recently been adopted by researchers and practitioners for supply chain systems where decision-makings are distributed (Caridi and Cavalieri, 2004).

Multi-agent based modelling and simulation techniques have originally been developed and applied in the Artificial Intelligence (AI) circle (Ferber, 1999). The technique is designed to handle simulation of distributed problem-solving systems where agents in the systems are autonomous or semi-autonomous in making decisions and responding to surrounding environments (Wooldridge, 2002). The modelling framework is generic in that it is applicable to a wide range of problems such as social science, biological science, physical science, computing, engineering, management, transportation etc. (Jennings and Wooldridge, 1998). The popularity of multi-agent based modelling is due to the fact that one could conveniently emulate a very complex system from a set of relatively simple agents (Parunak, 1997).

However, simulation models only model a pre-specified design to characterise the performance of that particular configuration. They cannot determine the most effective configuration from a set of potential configurations as optimisation does. Thus, if system dynamics is not a key issue, a static or basic "multi-period" model is appropriate and mathematical optimisation techniques can be applied.

Given the benefits and limitations of optimisation and simulation techniques, Hax and Candea (1984) suggested the two-stage approach which takes advantage of the strengths of both simulation- and optimisation-based approaches:

1. Use an optimisation model to generate a number of optimal solutions at the macrolevel, taking into account the most important cost components.
2. Use a simulation model to evaluate in more detail the solutions generated in the first phase.

Fu (2002) discussed approaches and algorithms of simulation optimisation tool in commercial software. In addition, he recommended that the desirable features in a good implementation of optimisation for commercial simulation software are generality, transparency to user, high dimensionality, and efficiency.

2.3.3 Types of Supply Chain Problems

As mentioned in the introduction, we can classify supply chain management analysis into supply chain execution (SCE) and supply chain design (SCD). A description of each of these is provided below with a review of literature in each area.

2.3.3.1 Supply Chain Operation and Execution

Supply chain operation and execution is the process of determining solutions to more tactical issues such as local inventory policies and deployment, manufacturing and service schedules, transportation plans, etc. In these instances, production and transportation data are usually assumed to vary according to a known probability distribution. The time period for the analysis typically spans days or weeks, and focuses on implementing detailed short term plans. Problems encountered most often by supply chain operation and execution are production planning, scheduling, etc.

Jung *et al.* (2004) proposed the use of deterministic planning and scheduling models which incorporate safety stock levels as a means of accommodating demand uncertainties in routine operation. The problem of determining the safety stock level to use to meet a desired level of customer satisfaction is addressed using a simulation based optimisation approach. They noted that the key limitation of the overall approach lies in the large computing times

required to address problems of increasing scope. They suggested that variance reduction techniques are viable means to improve computing time.

Bose and Pekny (2000) presented a model predictive framework which requires a forecasting model and an optimisation model working in tandem in a simulation environment for planning and scheduling problems and applied to a case study of a consumer goods supply chain. The model can help to suggest supply chain coordination structure – centralised control, decentralised control, or distributed control – under various circumstances.

Mohamed (1999) developed an integrated production and distribution model for a multi-national company operating in an environment of changing exchange rates. The objective is to minimise cost. The performance of the model is elicited through examples. He reported that the effect of exchange rate variation is felt worst when the capacity of the more profitable facility is restricted.

2.3.3.2 Supply Chain Network and Design

Supply chain design is the process of determining the supply chain infrastructure-the plants, distribution centres, transportation modes and lanes, production processes, etc. These studies are strategic in scope, use a time horizon of months or years, and typically assume little or no uncertainty in the data. Problems encountered most often by supply chain design are supply chain network configuration, supplier selection, location-allocation, etc.

Shapiro (2004) discussed new challenges for model practitioners and their clients due to the increasing number of applications of supply chain network optimisation models. He pointed out that current state-of-the-art models do not adequately describe decisions involving revenues, marketing campaigns, hedging against uncertainties, investment planning and other corporate financial decisions, and many other aspects of enterprise planning that interact with supply chain planning.

Ballou (2001) addressed unresolved issues in supply chain network design analysis and in the models used to support the analysis. He commented that systemwide logistics network configuring will likely be limited to relatively few echelons because of (1) information sharing problems among channel members, (2) the importance of the firm to upstream and down stream partners may not great enough to encourage cooperative location planning, and (3) the identification of the product over multiple echelons becomes blurred. Furthermore, even though a typical location model is probably adequate for most supply chain channel configuration

problems, there remain important development issues that can increase the accuracy, applicability, and scope of location analysis. Another important unresolved issue is which model solution approach is best. Too often “best” is argued on the basis of mathematical optimality when such issues as accuracy of problem representation, solution time, presentation features etc. can be of equal importance.

Harrison (2001) reviewed principles and methods of global supply chain design from two viewpoints. The first perspective is focused on the practitioner who is interested in an overview of the key concepts and applications of supply chain design within a global context. The second theme is to assess opportunities for research to extend supply chain design in useful directions, particularly into closely related areas such as supply chain operations and information systems. He concluded that the effect of duties, duty drawback, local content, taxes and exchange rates fluctuation are additional factors introduced by global supply chain networks. He suggested that a combination of optimisation with simulation techniques is a promising research direction to improve decision making for global supply chain design.

Cohen and Lee (1989) presented a normative model for resource deployment decisions in a global manufacturing and distribution network to support analysis of global manufacturing policies. The model also accounts for tariffs, duties, and transfer pricing. The results of applying the model to the analysis of global manufacturing strategies for personal computers were reported. In the application case, it demonstrated that profit rates, under the various alternatives, are not significantly different. The insensitivity of the profit rate illustrates the robustness of a global manufacturing network. It shows that the development of an international network of manufacturing and distribution resources provide the firm with the potential for increased flexibility. However, they noted that the mere presence of sites in different countries is not a sufficient condition for robustness. It was necessary for the model to generate new, optimal deployment decisions in responding to the changes associated with each of the different scenarios.

Cohen and Moon (1990) analysed the relationship between a firm’s manufacturing and distribution cost structure and characteristics of its supply chain strategy through Mixed-Integer Linear Programming (MILP) model. Their results indicate how scale economies, scope economies, and supply chain transportation costs impact network facility configuration, plant production charter, and distribution policy. The insight gained from their results was the demonstration that there can be dominant factors which determine aspects of the structure of the supply chain over a wide range of offsetting factors. Plant fixed costs (for a number of open

plants) and the cost of production complexity (for the number of products produced at a typical plant) are two examples of such dominant factors.

Tsiakis *et al.* (2001) considered the design of multi-product, multi-echelon supply chain networks. They proposed a model based on a detailed MILP with scenario analysis based on demand uncertainty. The decisions to be determined include the number, location, and capacity of warehouses and distribution centres to be set up, the transportation links that need to be established in the network, and the flows and production rates of materials. The results based on analysis of a particular case illustrated and validated the applicability of the approach.

2.4 Demand Forecasting

Demand is the main driving force of how supply chain planning and operations are conducted. However, notwithstanding the fact that supply chain planning requires some lead time to be executed, in most cases demand is not known in advance. Supply chain practitioners have to “guess” what demand will be in the future so that they can perform the supply chain planning now. In fact, the main concern of the supply chain is to look into the future, either in the short-term, medium-term, or long-term. In the same way, forecasting plays a central role in the operations function of a firm. All business planning is based on forecasts (Nahmias, 2005).

To highlight the importance of forecasting, let us take an example from a report on the front page of Metro on 20th December 2006 which was titled “£6 billion shopping spree in two days”. In the report a research director of TNS Worldwide, a market research company, warned supermarkets ‘*With the biggest Christmas week ever expected this year, they need to make sure they have enough in the warehouses to service both their Internet customers and those turning up to the stores.*’. Whereas Sainbury’s, a supermarket, said it had been preparing for Christmas since the summer as a spokeswoman added: ‘*We started planning at least six months ago to ensure we have enough stock.*’.

A carefully designed forecasting strategy can provide significant supply chain improvement even though there is no historical data to base forecasts on. An example is a case at Sport Obermeyer, a leading supplier in the US fashion ski apparel market, where the company could reduce significantly both stockout and markdown costs. Since the firm’s offerings are redesigned annually, historical data of previous sales are not relevant and the firm customarily based the forecasts on a consensus of members of the buying committee. However, in consensus forecasting, the dominant personalities in a group can influence other members’ forecasts; hence

a forecast obtained in this way might represent only the opinion of one person. By changing its forecasting procedure to individual forecasts supplied by members of the committee and other supply chain strategies, Sport Obermeyer experienced a cost reduction of about 5 percent of sales (Fisher *et al.*, 1994).

There are many methods to estimate the future demand of existing products and new products. The forecasting methods are divided into two categories: qualitative forecasting methods and quantitative forecasting methods. In practice, the demand forecast is frequently developed by a combination of quantitative and qualitative methods. This practice is in accordance with the views of forecasting academics who believe that better forecasts can be obtained by combining quantitative models with experts' opinions (Franses, 2004). The majority of practitioners are also reported routinely adjusting forecasts produced by software based on their judgment (Sanders and Manrodt, 2003).

2.4.1 Characteristics of Forecasts

There are some important characteristics of forecasts that all forecasts users should always bear in mind when they use or perform the forecasts. We list five of the most important points (Nahmias, 2005) below.

1. They are usually wrong. As strange as it may sound, this is probably the most ignored and the most significant property of almost all forecasting methods. Forecasts, once determined, are often treated as known information. Resource requirements and production schedules may require modifications if the forecast of demand proves to be inaccurate. The planning system should be sufficiently robust to be able to react to unanticipated forecast errors.
2. A good forecast is more than a single number. Given that forecasts are generally wrong, a good forecast also includes some measure of the anticipated forecast error. This could be in the form of a range, or an error measure such as the variance of the distribution of the forecast error.
3. Aggregate forecasts are more accurate. Recall from statistics that the variance of the average of a collection of independent identically distributed random variables is lower than the variance of each of the random variables; that is the variance of the sample mean is smaller than the population variance. This same phenomenon is true in forecasting as well. On a percentage basis, the error made in forecasting sales for an entire product line is generally less than the error made in forecasting sales for an individual item.

4. The longer the forecast horizon, the less accurate the forecast will be. This property is quite intuitive. One can predict tomorrow's value of the Dow Jones Industrial Average more accurately than next year's value.

5. Forecasts should not be used to the exclusion of known information. A particular technique may result in reasonably accurate forecasts in most circumstances. However, there may be information available concerning the future demand that is not presented in the past history of the series. For example, the company may be planning a special promotional sale for a particular item so that the demand will probably be higher than normal. This information must be manually factored into the forecast.

2.4.2 Forecasting Methods

Here we give an outline of major forecasting methods from both the qualitative and quantitative categories.

2.4.2.1 Quantitative forecasting methods

Quantitative forecasting methods utilise historical data to develop and verify patterns or relationships of the item to be forecasted versus time or other input variables. Generally, the patterns or relationships obtained will be described in terms of mathematical models as follows;

$$\text{Forecasts} = f(\text{input variables})$$

There are two types of quantitative forecasting models: time-series models and causal or explanatory models. A time-series model assumes that some pattern or combination of patterns is recurring over time. Thus, by identifying and extrapolating that pattern, forecasts for subsequent time periods can be developed. A time-series model describes forecasts as a function of time period. Hence, historical data series and the corresponding time periods are necessary to develop a time-series model. One advantage of time-series models is that the basic rules of accounting are oriented toward sequential time periods. However, this model also assumes that the underlying pattern can be identified solely on the basis of historical data series. This means that the model still generates the same forecast even if there are changes in other factors or actions that can affect the item to be forecasted. Examples of forecasting method using time-series models are Smoothing methods, Time Series Decomposition methods, and Autoregressive Integrated Moving Average methods.

Whilst a time-series model treats the system as a black box and makes no attempt to discover the underlying factors affecting the item to be forecasted, an explanatory model explicitly identifies such factors and describes forecasts as a function of those factors. The strength of an explanatory model is that once the factors or the input variables and the extent to which they affect the item to be forecasted are identified, a manager may be able to manipulate some variables to obtain a desirable level of forecasts. However, since these methods require information on several variables in addition to the variable that is being forecast, developing such methods can be very expensive or may be, in certain cases, impossible. Examples of explanatory model are Regression methods and Econometric models. Below we describe commonly used quantitative forecasting methods.

Smoothing Methods: Smoothing methods are a class of forecasting methods that use historical data to obtain a smoothed value for the series. The smoothed value is then extrapolated to become the forecast for the future value of the series. The smoothed value is obtained by averaging historical data, either as a simple moving average or exponentially weighted average. The basic notion inherent in exponential smoothing is that there is some underlying pattern in the values of the variables to be forecast and that the historical observations of each variable represent the underlying pattern as well as random fluctuations. The goal of these forecasting methods is to distinguish between the random fluctuations and the basic underlying pattern by smoothing (averaging) historical values.

Time series decomposition methods: Decomposition methods identify three separate components of the basic underlying pattern that characterize economic and business series. These are the trend, cycle, and seasonal factors. The trend factor, which represents the long-run behaviour of the data, can be increasing, decreasing, or unchanged. The cyclical factor represents the ups and downs caused by economic or industry-specific conditions. The cycle often follows the pattern of a wave, passing from a large to a small value and back again to a large value. The seasonal factor relates to periodic fluctuations of constant length and proportional depth that are caused by such things as temperature, rainfall, month of year, timing of holidays, and corporate policies. The distinction between seasonality and cyclicity is that seasonality repeats itself at fixed intervals such as a year, a month, or a week, whilst cyclical factors have a longer duration that varies from cycle to cycle.

Autoregressive Integrated Moving Average methods: Autoregressive Integrated Moving Average (ARIMA) methods are a class of time series methods that involve previous data values and previous values of errors. Autoregressive processes and moving-average processes are two models of time series with autocorrelation. Autocorrelation is the correlation among values of

observed data separated by a fixed number of periods. An autocorrelation of order one means the correlation among successive values of the observed data. The term integrated refers to differencing. Differencing is a means of eliminating trends and polynomial growth. This method exploits possible dependencies among values of the series from period to period. Accounting for these dependencies can often substantially improve forecasts.

ARIMA methods can deal with any data pattern. However, in most cases the manager must understand the method well enough to select the most appropriate specific model from the set available for that method. The Box-Jenkins methodology is a methodology for selecting among these models (Box and Jenkins, 1970). However, one difficulty with ARIMA models is that their handling requires quite a lot of algebraic manipulation. Furthermore, to obtain a reasonable estimate for the autocorrelation function, one must have a substantial history of observations. At least 72 data points are recommended (Nahmias, 2005).

Regression methods: As one type of explanatory model, in regression methods, a forecast will be presented as a function of a certain number of factors or variables that influence its outcome. The forecasts do not necessarily have to be time dependent. Developing an explanatory model facilitates a better understanding of the situation and allows experimentation with different combinations of inputs to study their effects on the forecasts. In this way, explanatory models can be geared toward intervention-influencing the future through decisions made today. Relationships between input variables can be linear or non-linear (e.g. polynomial, exponential) and the model can contain one input variable (simple regression) or multiple variables (multiple regression). Regression analysis uses the method of least squares to determine parameters (e.g. coefficients of the input variables) in the models.

Econometric models: Econometric models are systems of linear multiple regression equations, each including several interdependent variables. Whilst simple regression and multiple regression involve a single equation, econometric models can include any number of simultaneous regression equations. Therefore, one can also consider simple regression and multiple regression as special cases of econometric models. Below is an example of system of equations in an econometric model (Makridakis and Wheelwright, 1989).

$$\text{Sales} = f(\text{GNP, price, advertising})$$

$$\text{Production cost} = f(\text{number of units produced, inventories, labour costs, material costs})$$

$$\text{Selling expenses} = f(\text{advertising, other selling expenses})$$

$$\text{Advertising} = f(\text{sales})$$

$$\text{Price} = f(\text{production cost, selling expenses, administration overhead, profit})$$

The parameters of the equations are estimated in a simultaneous manner. The main advantage of econometric models lies in their ability to deal with interdependencies. However, complex econometric models are not always more accurate than simpler approaches such as time-series analysis or regression methods.

Quantitative forecasting methods are based on an assumption that trends or relationships exhibited in historical data will continue into the future. Hence, before using the results of quantitative forecasts, managers should consider whether this assumption is justifiable. If there is any information signifying changes in trend or relationship, such information should be taken into account and used such information to make an appropriate adjustment before delivering final forecasts.

2.4.2.2 Qualitative forecasting methods

Developing quantitative forecasts may not be possible or relevant when there is no historical data available or when the item to be forecasted is not a number (e.g. technology trend, fashion style). In such cases, qualitative forecasting methods or judgemental approaches to forecasting might be more appropriate. Below, we describe four major qualitative forecasting techniques.

The jury of executive opinion: The jury of executive opinion approach is one of the simplest and most widely used forecasting approaches available (see Dalrymple, 1987; Mentzer and Cox, 1984). The approach is to combine the opinions of experts to derive a forecast. Combining individual forecasts may be done in several ways. One is to have the individual responsible for preparing the forecast interview the executives and develop a forecast from the results of the interviews. Another is to require the executives to meet as a group and come to a consensus. When using this approach, a company generally brings together executives from sales, production, finance, purchasing, and administration so as to achieve broad coverage in experience and opinion. The advantage of the jury of executive opinion approach is that it provides forecasts quickly and easily. However, the main drawback of this approach is that the dominant personalities in a group can influence other members' forecasts. Thus the greatest weight will not necessarily be given to the assessment made by the member with the best information or the best ability to forecast.

The Delphi methods: In an effort to overcome the issue of personality in the jury of executive opinion approach, the Delphi method was developed at the Rand Corporation (Helmer and Rescher, 1959). The Delphi method is also based on soliciting the opinions of experts. However, the difference lies in the manner in which individual opinions are combined. The method requires a group of experts to express their opinions. The opinions are then made anonymous and compiled by a moderator. A summary of the results is returned to the experts, with special attention to those opinions that are significantly different from the group averages. The experts are asked if they wish to reconsider their original opinions in light of the group response. The process is repeated in hope that a consensus can be reached, or at least a reduced set of scenarios. The primary advantage of the Delphi method is that it provides a means of assessing individual opinion without usual concerns of personal interactions. The drawback is that it is highly sensitive to the care in the formulation of the questionnaire because there is no mechanism to resolve ambiguity through discussions.

Sales force composite methods: In the sales force composite approach, members of the sales force submit sales estimates of the products they will sell in the future. These estimates might be individual numbers or several numbers, such as pessimistic, most likely, and optimistic estimates. Sales managers would then be responsible for aggregating individual estimates to arrive at overall forecasts. Opinions of the sales force are used in this approach because the sales force has direct contact with consumers and is therefore in a good position to see changes in their demands. This technique is divided into three general categories (Hurwood *et al.*, 1978): the grass roots approach (collecting each salesperson's estimate), the sales management technique, and the distributor's approach. In certain situations, sales force composites may be inaccurate, for example if salespeople are overly optimistic or overly pessimistic, or when their compensation is based on meeting a sales target.

Market surveys and research methods: Whilst all qualitative forecasting techniques mentioned above rely on a group of experts to produce the forecasts, market surveys and research methods based the forecasts directly on survey results from a sample of the population (e.g. potential customers) whose behaviour and actions will determine future demands. The market research techniques are then used to gather survey information for organisation-specific needs. Market surveys can signal future trends and shifting customer preference patterns. To be effective, surveys and sampling plans must be carefully designed to guarantee that the resulting data are statistically unbiased and representative of the customer base (Tryfos, 1996; Bailar and Lanphier, 1978; Satin and Shastry, 1983).

Some major drawbacks of qualitative methods are judgmental biases, overconfidence in judgement, overconformity among group members, and inconsistency in judgement. Humans also have limited ability in identifying and quantifying underlying patterns and relationships in complex time series. For these reasons, if relevant historical data are available, the data should be used to develop quantitative forecasts with with an option of making adjustments by qualitative methods.

2.4.3 Evaluating Forecasts

Given the many forecasting methods available, how can one choose one forecasting method over the others? Chambers *et al.* (1971) elaborated a systematic approach in choosing the right forecasting technique. However, one of the main selection criteria is forecasting accuracy expressed in terms of residual errors as shown below.

$$e_i = F_i - D_i \quad (2.1)$$

where e_i is the forecast error in period i , F_i is the forecast made for period i , and D_i is the actual observed value during period i .

There are several indicators to represent forecasting accuracy. Here we introduce major forecasting accuracy indices that regularly appear in the literature and used by practitioners.

Mean Absolute Deviation (MAD) is one of the most commonly used forecasting accuracy measures. It is calculated by averaging absolute errors of each forecasting period. The absolute error is used to avoid negative errors and positive errors cancelling each other, in which average error will be underestimated. The formula of MAD is given below.

$$MAD = \frac{1}{n} \sum_{i=1}^n |e_i| \quad (2.2)$$

Mean Square Error (MSE) is another commonly used forecasting accuracy measure. It is the average of square errors. Since its form is similar to that of variance of a random sample, MSE is a preferable accuracy measure when further statistical analysis is to be performed (e.g. hypothesis testing). MSE penalises large forecasting errors more than small ones because errors are amplified by the power of two. However, it is more difficult to interpret than MAD because the unit of the errors is also squared. The formula of MSE is given below.

$$MSE = \frac{1}{n} \sum_{i=1}^n e_i^2 \quad (2.3)$$

Root Mean Square Error (RMSE) is the square root of MSE. It also has the property of giving a greater penalty to large errors like MSE but has the same units as the demand to be forecasted. The formula of MSE is as follows:

$$RSME = \sqrt{\frac{1}{n} \sum_{i=1}^n e_i^2} \quad (2.4)$$

Mean Absolute Percentage Error (MAPE) is not dependent on the magnitude of the values of demand. The absolute value of forecasting error is normalised by the corresponding actual demand. The average value of the normalised errors is then multiplied by 100 to convert into a percentage. MAPE is a relative accuracy measure, therefore it can be used to compare the forecasts of different products or of the same product but in different time periods. The formula of MAPE is given below.

$$MAPE = \left[\frac{1}{n} \sum_{i=1}^n \left| \frac{e_i}{D_i} \right| \right] \times 100 \quad (2.5)$$

In some forecasting methods, such as regression and ARIMA, the residual errors of the model also need to be random, constant, and normally distributed. In other methods (e.g. smoothing methods, time series decomposition methods), there is no restriction about the residual errors. However, it is desirable that they be random, constant, and normally distributed. In addition to forecasting accuracy measures above, calculating the mean error ($ME = \frac{1}{n} \sum_{i=1}^n e_i$) can also reveal possible forecasting bias.

In addition to standard methods mentioned above, several advanced forecasting techniques have been introduced in the literature. For example, Levis and Papageorgiou (2005) presented an optimisation-based approach for demand forecasting through support vector regression (SVR) analysis. They proposed a three-step algorithm using both nonlinear programming and linear programming to determine the regression function while the final step employs a recursive methodology to perform demand forecasting. The applicability of the proposed method was then validated through three demand forecasting examples which all achieved more than 93% prediction accuracy.

2.5 Supply Chain Complexity and Complexity Measurement

That supply chain systems nowadays are increasingly complex is generally agreed by most researchers and practitioners in supply chain management. In this section, we present some interesting literature addressing supply chain complexity.

Choi and Krause (2006) studied complexity in a supply base, the portion of a supply network that is actively managed by a buying company. They conceptualise supply base complexity in three dimensions: (1) the number of suppliers in the supply base, (2) the degree of differentiation among these suppliers, and (3) the level of inter-relationships among the suppliers. They discover that although a reduction in complexity may lead to lower transaction costs and increased supplier responsiveness, in certain circumstance it may also increase supply risk and reduce supplier innovation.

Wu *et al.* (2007) attempted to address the relationship between costs and the complexity index developed by Frizelle and Woodcock (1995). They carried out measurements on two types of supplier-customer systems in the UK; the make-to-stock system with low product variety but high volume and the make-to-order system with high variety but low volume. They found that the operational complexity is associated with the operational costs of running a supply chain. Moreover, the costs can be apportioned between those associated with the structure of the supply chain and those generated by departures from what was planned.

However, the methodology of quantifying complexity appropriately is still a conundrum. In this section we present some attempts in measuring complexity with a discussion of advantages and problems of such methods.

2.5.1 The Information-theoretic Measure of Complexity

Several papers on this method have been proposed by Frizelle (1995), Calinescu *et al.* (2000), and Sivadasan *et al.* (2002). They propose two fundamental complexity measures; structural complexity and operational complexity. This method defines complexity in terms of the expected amount of information required to describe the state of the system (structural complexity) and the extra amount of information required when the system deviates from expectation (operational complexity). Quantification of information is done by means of “entropy” which was first introduced by Shannon (1948). Calculation of structural complexity and operational complexity are given below.

$$H_s = -\sum_{i=1}^M \sum_{j=1}^{S_i} p_{ij} \log_2 p_{ij}$$

$$H_D = -P \log_2 P - (1-P) \log_2 (1-P) - (1-P) \sum_{i=1}^M \sum_{j \in NS_i} p_{ij} \log_2 p_{ij}$$

where H_s is quantity of structural complexity

H_D is quantity of operational complexity

M is number of resources

S_i is possible states of the i -th resource

p_{ij} is the probability of resource i being in state s_{ij}

P is probability that the system is in an *In Control & Planned* state

NS_i is an *unplanned* state or non-scheduled state

An example of a calculation for structural complexity is as follows:

Imagine a production system composed of three machines that operate sequentially. The system produces only four goods and every machine takes part in the production of each of the four products. This means that each production unit can be in one of four states according to the product it is working on and without taking into account of idle time and breakdowns. Therefore, there are four states ($S_i=4$) for each machine. Assume that there is stochastic independence among the states and that all the states are equally likely. Thus, since $p_{ij} = 1/4$ we obtain the following result.

$$H_s = -\sum_{i=1}^M \sum_{j=1}^{S_i} p_{ij} \log_2 (p_{ij}) = -\sum_{i=1}^3 \sum_{j=1}^4 \frac{1}{4} \cdot \log_2 \left(\frac{1}{4} \right) = -\sum_{i=1}^3 \sum_{j=1}^4 \frac{1}{4} \cdot (-2) = -3 \cdot 4 \cdot \frac{1}{4} \cdot (-2) = 6$$

Operational complexity can be calculated in similar manner, though it is more complicated.

Wu *et al.* (2002) studied a simulation on supply chain operational complexity in the manufacturing industry based on the information theoretic view. The simulation was carried out using a case from a major UK company and its suppliers. It mimics the information supplied by the customer, models the production of the supplier, and reproduces the planning and scheduling of the production. Complexity indices are calculated. They found that the simulation results are consistent with those found when the real data were used in the model. The simulation also supports a tentative conclusion made based on the real data that communications between customer and supplier can improve the performance of the supply chain (e.g. reduce stock out rate).

Yu and Efstathiou (2002) presented an entropic complexity measure to indicate the effect of modification in the manufacturing network. It provides a quantitative method to evaluate the performance of system layout design in terms of structural complexity and production rate. The framework was proposed to analyze the effect of different manufacturing network configurations. Their new methodology enables calculation of the manufacturing network complexity using a step by step procedure.

2.5.2 The Chaos-theoretic Measure of Complexity

The chaos-theoretic measurement concerns only the dynamic part of complexity, which is equivalent to operational complexity. This method is proposed by Deshmukh (1993) as a measurement of complexity in manufacturing systems in terms of the rate at which a system loses information per period. Chaos theory investigates the properties of systems that can be defined according to equations of the form:

$$x_{n+1} = f(x_n)$$

Such systems may show extreme sensitivity to the initial value of x i.e. x_0 . However, time-plots of the state of the system may show its behaviour as being approximately cyclical, i.e. orbiting around an attractor. If the system is diverging, then the distance between pairs of points should be increasing exponentially according to:

$$|x_{n+1} - x_n| = A \cdot e^{\lambda_n}, \lambda > 0$$

where $A = |x_1 - x_0|$. Rearranging, we obtain:

$$\lambda = \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{k=0}^{n-1} \ln |f'(x_k)|$$

A value of λ may be calculated for each period of the system. If the logarithm is taken to the base 2, this gives the number of bits of information that are being lost to the system with each period.

Efstathiou *et al.* (2001) presented some preliminary findings on the relationships between complexity measures arising from information theory and chaos theory. The authors showed that an application to a case study data on supply chain stock demonstrates period doubling and orbital behaviour. They claimed that both approaches can be used to derive the feasible horizon of a production schedule defined as a ratio of the expected amount of information contained in

one issue of a schedule to the actual information generated in non-scheduled states. The authors suggest that a comparison of feasible horizon with the actual horizon would yield useful insights on the ability of a system to stick to its schedule.

Both the information-theoretic approach and chaos-theoretic approach rigorously define complexity from a theoretical point of view. However, practically, measuring complexity with such methods requires a large amount of time and data to adequately observe all the possible states. Furthermore, defining states in the system is also a subjective choice. Different evaluators may come up with different sets of states resulting in different resolution in measuring complexity of the same system. Finally, complexity by these definitions does not always mean “good” or “bad”. Some organisations may intentionally maintain high levels of complexity in their systems without any additional cost or losing competitive advantage.

2.5.3 The Fitness Measure of Complexity

Fitness measures view complexity in terms of degrees of interdependence. The concept of the fitness measure of complexity stems from biological systems and has been recently adopted in the field of economic analysis. One of the most successful formulations is Stuart Kauffman’s NK model. Kauffman (1993) characterises the fitness landscapes essentially with two structural variables:

- N, which is the number of elements that characterises the entity;
- K, which is the number of elements of N with which a certain attribute interacts.

In order to describe the NK model it is therefore necessary, as a first step, to define a system composed of N elements that can assume several states (e.g. a production process composed of N machines). Such elements can be characterised by different degrees of interdependence. Thus, the variable K expresses the average number of other elements, which every element is interdependent with. Every possible combination of states of the single elements is called a configuration. Finally, it is necessary to define a measure of performance of the system that is called fitness. The set of values of fitness constitutes the fitness landscape. For example, in a production system, we can consider the performance as a combination of the time required in order to complete a production process (a smaller process time means a better performance) and of the cheapness in the use of the resources.

In absence of interdependencies (that is for $K = 0$), we have a landscape characterised by only one global optimum point. This can be explained by the fact that, if the contribution of each single element to the overall fitness is independent of the others, then there is a characteristic

which is independent of the one of the other elements. This amounts to say that every element can improve its own contribution without interfering with the characteristics of the others. The lack of interdependencies in the fitness landscape implies a situation in which we have only one global optimum and every element improves the overall fitness through local optimisation.

When the level of interdependence is high, a change of a single action can lead to a decrease in the fitness although a simultaneous change of a greater set of actions can improve the performance. This means the improvement of only one component may decrease the overall performance of the entire system. For example, adopting the system known as *just in time* for the management of inventories can decrease the efficiency of the organisation in the case in which other contemporary changes in the production and management systems are not made.

Obviously systems with greater levels of interdependency tend to require more resources to coordinate for global optimality. This view of complexity is usual in evaluating the operational behaviour of a system. However, it is not clear how to evaluate the cost and benefit of interdependency.

Gino (2002) discussed complexity measures in decomposable structures. Decomposable structures are those in which components are minimally interactive and interaction can be handled in an additive or perhaps linear way. Two kinds of complexity measures are considered; entropy measures based on information theory and fitness measures based on the number of interdependent relationships within a system. The author compared the two groups of methods and concluded that entropy is an additive measure with perfect decomposability; the entropy of the system is the sum of the entropy of each unit. However, the entropy measure focuses only on the information complexity of a system without taking account of the interdependencies existing among the components. On the other hand, fitness measures explicitly consider this feature and can address systems with different degrees of decomposability.

2.5.4 Complexity as degrees of connectivity

Another possible complexity metric is the degree of connectivity. By this view, a system is composed of several elements, and highly connected systems are more complex than less connected systems with the same number of components. This view of complexity can explicitly reveal the cost of complexity. The more highly connected the system or network, the greater the flexibility of the system. However a highly connected system requires a high level of

coordination, a significant amount of information for decision-making, and finally is costly to manage under high complexity.

Venkatasubramanian *et al.* (2004) define excess connectivity of networks as redundancy. They introduce the measure *redundancy coefficient*, which is defined as the excess number of edges over an MST (minimum spanning tree) normalised by the maximum possible departure from an MST. They categorise redundancy into two types; functional redundancy (if a redundant edge is added between vertices that are not directly connected) and structural redundancy (if otherwise), respectively.

Each type of complexity mentioned above defines and measures complexity from different points of view. The information-theoretic method, the chaos-theoretic method, and the fitness method are suitable for measuring complexity of process or operations, for example, operations of production process or distribution system. The information-theoretic method measures the level of information required, whilst the chaos-theoretic method measures the rate at which a system loses information per period and the fitness method measures level of interdependence. Depending on information available and the characteristics of the system, we may be able to calculate one or more types of complexity measure. The degrees of connectivity method measures the level of connection within the system, hence it is suitable for measuring complexity of infrastructure or network.

2.6 Supply Chain Robustness

The design and establishment of supply chain networks is a strategic decision whose effect will last for several years, during which the business environment may change. Furthermore, some undesirable events (e.g. earthquakes, sabotage, accidents) that can severely disrupt the supply chain may happen at some point of time during the lifespan of the supply chain network even though they may happen infrequently and with very low probabilities. For example, in September 1999, an earthquake in Taiwan delayed shipments of computer components to the United States by weeks and, in some cases, by months. In March 2000, a Philips facility in Albuquerque, New Mexico, went up in flames (Lee, 2004). This can significantly damage a supply chain's performance if they are the only sources of particular components and if final products cannot be produced without them. Supply chain robustness can be defined as the extent to which the supply chain is able to carry out its functions despite some damage done to it, such as the removal of some of the components in the supply chain network.

From a network point of view, Venkatasubramanian *et al.* (2004) presented a theory of self-organisation by evolutionary adaptation in which they showed how the structure and organisation of a network is related to survival, or in general the performance, objectives of the system. The survival of a network depends on efficiency (short-term survival), robustness (long-term survival), and cost. Such networks can be found in a variety of applications including supply chain networks. Using a graph theoretical case study and optimisation by Genetic Algorithms simulation, they showed that when efficiency is paramount the “star” topology emerges and when robustness is important the “circle” topology is found. When efficiency and robustness requirements are both important to varying degrees, other classes of networks such as the “hub” emerge.

Regarding supply chain performance, Lee (2004) argued that *efficient supply chains* often become uncompetitive because they do not adapt to changes in the structures of markets. Such supply chains usually focus only on high speed and low cost which will work brilliantly under planned conditions but when unexpected changes in demand or supply occur, they are unable to respond. He suggested that only supply chains that are agile, adaptable, and aligned provide companies with sustainable competitive advantage. Agility is the ability to respond to short-term changes in demand or supply quickly and to handle external disruptions smoothly. Adaptability is the ability to adjust a supply chain’s design to meet structural shifts in markets; modify supply networks to strategies, products, and technologies. Alignment means the alignment of interests of all players in the supply chains to create incentives for better overall performance. Finally, he concluded that supply chain efficiency is necessary, but it is not enough. Only those companies that build agile, adaptable, and aligned supply chains will stay ahead of their competition.

The current world situation is another imperative to the requirement of robust supply chains. Sheffi (2001) asserted that on the morning of September 11th, 2001, the United States and the Western world entered into a new era - one in which large scale terrorist acts are to be expected. The impacts of the new era will challenge supply chain managers to adjust relations with suppliers and customers, contend with transportation difficulties and amend inventory management strategies. He looks at the twin corporate challenges of (i) preparing to deal with the aftermath of terrorist attacks and (ii) operating under heightened security. The first challenge involves establishing certain operational redundancies. The second means less reliable lead times and less certain demand scenarios. He emphasises that companies should organise to meet these challenges efficiently and suggests a public-private partnership for supply chain management.

Supply chain robustness is increasingly gaining attention in the supply chain management research community as it is a crucial strategy for avoiding supply chain breakdowns (Chopra and Sodhi, 2004), mitigating supply chain disruptions (Christopher, 2006), and creating resilient supply chains (Christopher, 2004).

2.7 Literature Review Summary and Discussion

A large body of supply chain management, complexity measurement, and robustness related work has been reviewed in this chapter. Supply chain management virtually encompasses a firm's activities from the strategic level to the operational level. Amongst these activities, strategic level activities bear the greatest impact on supply chain performance, in terms of both time-horizon and return on investment.

Unfortunately, there is no one strategy that fits all circumstances in supply chain management. Though highly efficient and highly successful in the US, Wal-Mart's supply chain has struggled desperately in Japan due to differences in supply chain environment and customer value. Typically, Japanese customers do not place high value on low price. They value high product quality and services. Thus, Wal-Mart's "Every Day Low Prices (EDLP)" strategy does not work well in Japan (Rowley, 2005).

When the environment changes or the environment becomes uncertain, a previously highly effective supply chain may become a very ineffective one. The example above is only one out of a myriad of similar problems faced by practitioners around the globe. Wal-Mart is forced to confront two options; rely heavily on marketing to change Japanese customers' attitudes, or re-organise its supply chain to match customer value.

Re-organising a supply chain is always a painful task. It might be useful to design a supply chain such that risk of environmental change is taken into consideration at the design stage so that the supply chain is, to a certain extent, robust against the change.

Whilst a large number of supply chain design papers discuss efficiency and responsiveness in the supply chain, papers on quantifying robustness appeared more in the generic network literature. Coupling the concepts of efficiency and robustness in strategic supply chain network design and management is a very interesting area to explore.

Research in complexity measurement has recently appeared in the field of manufacturing and supply chain management. Though its cost and benefit are not always obvious, linking supply chain complexity measurement to a trade-off between efficiency and robustness is certainly of great potential to unveil insights in supply chain performance that have not been made so far. This is the objective of later chapters. We will explore relationships of supply chain efficiency, robustness, and complexity through mathematical models, simulations, and case studies. We will also characterise and investigate supply chain networks under graph theoretic formalisms to observe their consistency with the power law found in many other natural and man-made networks (Albert and Barabási, 2002).

Chapter 3

Logistical Network Design

Logistical network design is one of the earliest strategic decisions in supply chain management that supply chain managers have to make. Practitioners and researchers typically focus on optimising efficiency and/or responsiveness of logistical networks. This chapter presents a logistical network design framework with robustness and complexity considerations. We define robustness, complexity, and normalised efficiency of a logistical network. A mathematical model is constructed based on our conceptual framework and applied to a hypothetical case study with varying robustness requirements. Since all functions are linear and some variables can only be integer, the mathematical model is formulated as a Mixed-Integer Linear Programming (MILP) problem which is a similar approach to Tsiakis *et al.* (2001). Furthermore, we introduce a graph-theoretic view of the logistical network and present its topological properties such as average path length, clustering coefficient, and degree distribution. We also present the relationships of robustness versus normalised efficiency and complexity.

3.1 Logistical Network Design Concept

The main concept of logistical network design in this chapter is to explicitly incorporate robustness into the network design specification. This section describes the development of a logistical network design model and performance measurements of the networks. A case study and results of the models will be presented including topological properties of logistical networks. The most widely used method for supply chain network design is Mixed-Integer Linear Programming (MILP). Here, we consider the design of a multi-echelon logistical network as follows:

Objectives

A typical objective in the design of logistical networks in supply chains is to maximise profit or minimise cost whilst satisfying all established constraints. Other possible objectives include, but are not limited to, maximisation of market share, return on investment, asset utilisation, or minimisation of inventory, customer delays etc. In general, these objectives can be collectively termed as “efficiency”. However, as already mentioned, we believe that efficiency should not be the only performance measure of logistical networks that supply chain managers should be concerned about. As time passes, environmental pressures and/or sudden changes may severely impact the system such that the entire the supply chain may become defunct, not to mention efficiency being compromised. This long-term survival aspect should naturally be taken into account. In what follows, we demonstrate different logistical network structures based on various trade-offs between efficiency and robustness including network complexity measurement. First we provide definitions of efficiency, robustness, and complexity of logistical networks, and graph theoretic formalisms, and then we develop a mathematical model, and demonstrate and interpret the results.

Network Efficiency

Efficiency is a measure of the effectiveness of the system’s configuration to accomplish its function. We propose efficiency of logistical network to be measured by conventional supply chain optimisation objectives such as costs, customer delays etc. (Chopra and Meindl, 2004). In this paper, we use cost as our measure of absolute efficiency. For example, distribution costs consist of fixed set up costs, handling costs, transportation costs etc. The most efficient logistical network is the network that requires the lowest cost. The normalised efficiency (η_s) of the logistical network is the ratio of its absolute efficiency to that of the most efficient network for the same problem setting:

$$\eta_s = \frac{Eff^*}{Eff_s}$$

Eff_s = Absolute efficiency (cost) of the logistical network of interest

Eff^* = Absolute efficiency (cost) of the most efficient logistical network

Network Robustness

Robustness is the extent to which the system is able to carry out its functions despite some damage done to it, such as the removal of some of the nodes and/or edges in the network. The

objectives of design for efficiency and design for robustness are often in conflict. Whilst one tries to improve robustness, efficiency is generally compromised.

Based on the general definition of robustness, here we define our own measure of logistical network robustness as the ratio of the extent to which the network is able to fulfil demand when deleting component j (β_j) from the network and total demand in the network. In other words, it is a normalised demand-weighted effective accessibility of a network.

$$\beta_j = D'_j / D$$

where, D'_j = amount of demand that can be fulfilled when deleting j and D is the total demand in the network.

The average-case robustness of a network is then defined as the average of the robustness values computed over all the components:

$$\beta_{avg} = \frac{1}{N} \sum_{j=1}^N \beta_j$$

The worst-case robustness is the minimum of the robustness values computed over all the components:

$$\beta_{worst} = \min_j \beta_j$$

Here, we shall constrain the worst-case robustness in our model. However, we shall compute the average case robustness of our results for illustration purposes. In addition, unless clearly specified as otherwise (e.g. average case robustness), hereafter, the robustness shall mean the worst-case robustness. The terms “minimum robustness” and “the worst-case robustness” will also be used interchangeably in this chapter. According to this definition, appropriate excess links (redundancy) will increase the robustness of a network (more channels to fulfil demand). However, sometimes robustness can be improved without adding more edges. The overall robustness of a network depends on two features: its *topological robustness* and *redundancy robustness* (Venkatasubramanian *et al.*, 2004).

Network Complexity

Information theoretic or entropic views of complexity measurement seem to dominate the majority of the literature we have surveyed. However, the mechanism of the entropic approach which requires state probabilities and the characteristics of logistical network topology are quite different. In contrast to these methods, we shall develop another characterisation by adapting a redundancy view of complexity. By our definition below, a highly linked network is more complex than a less linked network. Logistical network complexity is then defined as the ratio of actual number of network edges to the number of minimum spanning tree edges (the minimum number of links required to connect all components of the network) of the smallest network of the same problem setting.

$$C = \frac{E_n}{E_{MST}^*}$$

E_n = actual number of network edges

E_{MST}^* = number minimum spanning tree edges of the smallest network

By this definition, the more highly linked the network, the more complexity it possesses. This complexity definition can easily be related to cost as well.

Graph Theoretic Formalism

We introduce a graph theoretic framework into the analysis of the logistical network topology. Components in the logistical network (plant, warehouse, customer etc.) are referred to as nodes or vertices. Links between components or vertices are called edges. The robustness of the logistical network depends on the ability of each vertex to communicate or interact, directly or indirectly, with all other vertices in the network. The interactions can be through the flow of material, goods, or information. As long as a vertex is connected to another vertex, it can communicate with it and pass on the flow to others which are part of the overall connectivity. Below are some definitions of terms in the graph theoretic framework (Deo, 1974).

Vertex degree or degree (k) of a vertex is the number of edges attached to it. We define $n(k)$ as the number of vertices with degree k in the graph and $p(k)$ is defined as the probability that any vertex will have degree k . The average degree of vertex is the average degree over all vertices in the graph and is defined as follows:

$$\langle k \rangle = \text{Average degree of vertex} = \frac{\sum_{i=1}^N k_i}{N}$$

The distance $d(i, j)$ between vertices i and j of a graph is the number of edges along the shortest path between them. The average path length is then the average distance between any pair of vertices:

$$\text{Average Path Length (APL)} = \frac{\sum_{i,j} d(i, j)}{N(N-1)/2}, \quad 1 \leq i, j \leq N$$

The tendency of a graph to cluster is quantified by the clustering coefficient (Watts and Strogatz, 1998). Consider a vertex i in a graph having k_i edges which connect it to k_i other vertices. If the nearest neighbours are completely connected (i.e. there are direct edges from any vertex to every other vertices), there would be $k_i(k_i - 1)/2$ edges between them. The ratio between the number E_i of edges that actually exist between these k_i vertices and the total number $k_i(k_i - 1)/2$ gives the value of the clustering coefficient of vertex i ,

$$C_i = \frac{2E_i}{k_i(k_i - 1)}$$

The clustering coefficient of the whole graph is the average of all individual C_i .

A graph or network is said to be connected when one can go from any vertex to any other vertex by following the edges. For a network with N vertices, the minimum number of edges required to obtain a connected graph is $E_{min} = N - 1$ (Clustering coefficient = 0). This reduced set which defines the skeleton of a network is typically known as the *Minimum Spanning Tree* (MST).

Application of the design concept

As a demonstration of our concept, we introduce a three-echelon supply chain network design problem with one manufacturing facility. Manufactured products are distributed from the manufacturing plant to customer zones via warehouses. For simplicity, only one product type (aggregated demand forecast and no handling cost differentiation between products) is considered in our sample application and the manufacturing plant location is fixed. The objective is to minimise overall cost; a robustness requirement is introduced by adding the following features into the mathematical model.

1. There is a requirement that, in the worst case, the network must be able to fulfil a minimum fraction of demand.

2. A primary link between a customer and a warehouse is established when a customer's demand is served by a warehouse. However, it is also possible to establish a secondary link for standby between a customer and another warehouse with some small extra cost. When the primary warehouse fails to function, the warehouse in secondary link will fulfil the customer's demand instead. Secondary links are only possible in logistical networks where there are two or more warehouses.

3.2 Mathematical Model

Design formulation

The following decisions are considered to optimise our objective function.

- Number of warehouses
- Warehouse location
- Assignment of warehouses to serve customers (primary and secondary)

Network structure constraints

There are n customer zones. Warehouses can be established in any customer zone and the number of warehouses is limited to h

$$\sum_{i=1}^n W_i \leq h \quad (3.1)$$

where W_i is a binary variable such that

$$W_i = \begin{cases} 1 & \text{if a warehouse is to be established in zone } i \\ 0 & \text{otherwise} \end{cases}$$

Each customer will be assigned to a warehouse. A link between customer zone j and a warehouse in zone i is possible only if there is a warehouse in zone i

$$X_{ij} \leq W_i, \quad \forall i \quad (3.2)$$

where X_{ij} is a binary variable such that

$$X_{ij} = \begin{cases} 1 & \text{if customer zone } j \text{ is assigned to a warehouse in zone } i \\ 0 & \text{otherwise} \end{cases}$$

In addition, each customer zone will be served by exactly one primary warehouse (single sourcing requirement from primary link).

$$\sum_{i=1}^n X_{ij} = 1, \quad \forall j \quad (3.3)$$

Demand Constraints

The quantity of product distributed from a warehouse in zone i to customer zone j is Q_{ij} .

Product from a warehouse in zone i can be distributed to customer zone j only when a link between them exists.

$$Q_{ij} \leq M \cdot X_{ij}, \quad \forall i, j \quad (3.4)$$

where M is an appropriately large number.

The customers in zone j require product of quantity D_j , which must be fulfilled under normal circumstances.

$$\sum_{i=1}^n Q_{ij} = D_j, \quad \forall j \quad (3.5)$$

Material balance requirements

The quantity of product distributed from plant to warehouse in zone i is P_i . Product from the plant can be distributed to warehouse i only when a warehouse is established in zone i .

$$P_i \leq M \cdot W_i, \quad \forall i \quad (3.6)$$

again M is an appropriately large number.

The product quantity required by warehouse i must be equal to the demand of all customers assigned to warehouse i .

$$\sum_{j=1}^n Q_{ij} = P_i, \quad \forall i \quad (3.7)$$

Secondary links and robustness requirement

It is possible to establish a secondary link to each customer zone, however, there cannot be a secondary link from the warehouse that already serves as a primary source to that customer zone.

$$X_{ij} + XS_{ij} \leq 1, \forall i, j \quad (3.8)$$

where XS_{ij} is a binary variable such that

$$XS_{ij} = \begin{cases} 1 & \text{if a secondary link is established between customer zone } j \text{ and warehouse } i \\ 0 & \text{otherwise} \end{cases}$$

A secondary link can be established from any zone only if a warehouse exists in that zone.

$$XS_{ij} \leq W_i, \quad \forall i, j \quad (3.9)$$

At most one secondary link is allowed for any customer zone in the supply chain network.

$$\sum_{i=1}^n XS_{ij} \leq 1, \quad \forall j \quad (3.10)$$

If no secondary link exists to a customer zone, the demand of that customer zone cannot be fulfilled if its primary warehouse fails. Now we introduce a binary variable which defines this relationship.

$$XF_{ij} \geq X_{ij} - \sum_{i'=1, i' \neq i}^n XS_{i'j}, \quad \forall i, j \quad (3.11)$$

where XF_{ij} is a binary variable such that

$$XF_{ij} = \begin{cases} 1 & \text{if demand of customer zone } j \text{ cannot be served if its primary warehouse } i \text{ fails} \\ 0 & \text{otherwise} \end{cases}$$

Finally, we specify a requirement of minimum robustness (in the worst case, overall fulfilment of customer demand must not fall below this level) of the logistical network.

$$1 - \frac{\sum_{j=1}^n XF_{ij} \cdot D_j}{\sum_{j=1}^n D_j} \geq \beta, \quad \forall i \quad (3.12)$$

where β is a minimum robustness and $0 \leq \beta \leq 1$. The numerator of the second term on the left hand side of the constraint is the fraction of demand that cannot be fulfilled when a warehouse i fails to function, whilst the denominator is the total demand in the network. Hence, the left hand side of the constraint represents the fraction of demand that still can be fulfilled when a particular warehouse i fails.

Non-negativity constraints

$$Q_{ij} \geq 0, \quad \forall i, j \quad (3.13)$$

$$P_i \geq 0, \quad \forall i \quad (3.14)$$

The objective function of our model is cost; below we define each component of total logistical network cost.

Fixed infrastructure cost

The infrastructure cost considered in this basic model is related to the establishment of a warehouse at a particular zone. We assume that the establishment of any warehouse incurs an equal cost C^w per period (e.g. amortised fixed set up cost throughout the life time). The number of warehouses to be established depends on the solution of a model, and the total fixed infrastructure cost per period is:

$$\sum_{i=1}^n C^w \cdot W_i$$

Material handling cost

The material handling cost is assumed to be a linear function of the total throughput incurred at the warehouses. In our model, we further assume that handling costs per unit product, C^h , at all warehouses for all customers are the same. The total handling cost then becomes:

$$\sum_{i=1}^n \sum_{j=1}^n C^h \cdot Q_{ij}$$

Transportation cost

A transportation cost is incurred when product is distributed from one node to another node. For this model, it is from the plant to warehouses and then from warehouses to customers.

Transportation cost is a function of unit rate, distance, and quantity of product. Therefore, the cost of transportation from the plant to warehouses is

$$\sum_{i=1}^n C^t \cdot P_i \cdot LP_i$$

where C^t is a transportation unit cost rate and LP_i is distance between the plant to each warehouse in zone i , if exist. Notice that if no warehouse exists in zone i , P_i will be zero and $P_i \cdot LP_i$ will be zero. Thus, no transportation cost is incurred for non-existent warehouses. The cost of transportation from warehouses to customers is:

$$\sum_{i=1}^n \sum_{j=1}^n C^t \cdot Q_{ij} \cdot L_{ij}$$

where L_{ij} is the distance between a warehouse in zone i and customer zone j . Notice again that if there is no link between a warehouse in zone i and customer zone j , Q_{ij} will be zero and its multiplication will be zero. Again, no transportation cost is incurred for non-existent links between warehouses and customers.

Cost of secondary links

An additional cost is incurred when a secondary link is established for standby. We assume that this additional cost is proportional to demand and the distance between the secondary warehouse and the customer zone. It will relate to the management and contractual costs of establishing the link as well as actual logistics costs on the occasions that the links is used. This additional cost can then be defined as:

$$\alpha \sum_{i=1}^n \sum_{j=1}^n D_j \cdot L_{ij} \cdot XS_{ij}$$

where α represents a fraction of the full transportation cost rate and $0 \leq \alpha \leq C^t$

Hence, the objective function is:

$$\min \sum_{i=1}^n C^w \cdot W_i + \sum_{i=1}^n \sum_{j=1}^n C^h \cdot Q_{ij} + \sum_{i=1}^n C^t \cdot P_i \cdot LP_i + \sum_{i=1}^n \sum_{j=1}^n C^t \cdot Q_{ij} \cdot L_{ij} + \alpha \sum_{i=1}^n \sum_{j=1}^n D_j \cdot L_{ij} \cdot XS_{ij} \quad (3.15)$$

Notice that there is no inventory cost in the objective function. This is reasonable if we view demand as deterministic (e.g. locked by a long-term contract) or if we use warehouses as transfer or cross-docking points rather than storage points. In both cases, no significant inventory is required at the warehouses and all inventories are held centralised at the plant. Hence, there is no difference in inventory cost no matter how many warehouses are established. For simplicity, we therefore do not include inventory cost in the objective function since it will not affect the solution.

3.3 A Case Study and Results

In our case study we consider customers located in the following customer zones denoted by number 1 to 36.

Customer Zone

1	2	3	4	5	6
7	8	9	10	11	12
13	14		16	17	18
19	20	21	22	23	24
25	26	27	28	29	30
31	32	33	34	35	36

Figure 3.1: Customer zone geography

The plant location is fixed at zone 15. The area of each zone width is 10 km vertically and 10 km horizontally. The distance between zones is supposed to be the direct Euclidean distance between centres of each zone. Monthly product demand for each zone is generated randomly using a uniform distribution ranging from 0 to 1,000 ton as follows:

393.5	45.5	744.9	813.6	919.6	559.0
158.6	153.4	618.8	882.3	78.7	706.5
973.4	452.5	358.1	440.6	593.0	321.3
624.4	301.8	321.0	821.0	913.8	399.2
542.1	374.9	447.0	30.4	751.6	701.8
202.4	793.3	829.3	942.0	899.9	303.1

Figure 3.2: Customer demand distribution (corresponding to customer zones in Figure 3.1)

The unit transportation cost rate is constant at £2 per ton per km, and the handling cost rate is £0.5 per ton. The maximum number of warehouses is set at 3. The fixed set up cost for each warehouse is £4,000/month. We set $\alpha = £0.05$ and solve six different scenarios by varying β (the worst-case robustness requirement) as 0.0, 0.1, 0.3, 0.5, 0.7, and 0.9 respectively.

Model solving and results

The case study can be formulated and solved as a Mixed-Integer Linear Programming (MILP) problem. When there is no robustness requirement ($\beta = 0$), the most efficient network topology turns out to be a 'star' network with only one warehouse to be established at the same location as the plant since this network required the lowest cost, or the highest efficiency. Since it requires only one warehouse, the star network is the most compact network in this problem

setting. It is also well known that a star network forms a minimum spanning tree. The optimised network is illustrated in Figure 3.3.

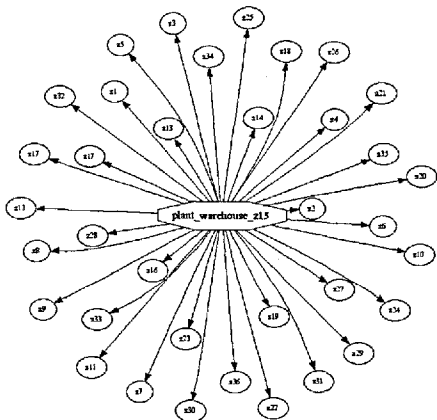


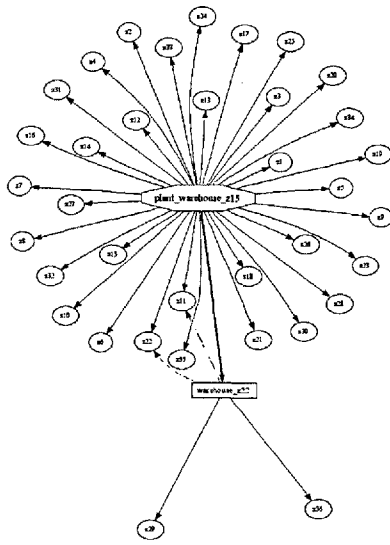
Figure 3.3: Logistical network structure (no robustness requirement)

As we can see from Table 3.1, the worst case robustness of the logistical network in Figure 3.3 is zero. That is if the only warehouse fails to function, the whole supply chain becomes defunct. Performance measures of logistical networks with robustness requirements are also listed in the Table 3.1. The model can produce logistical networks with the worst-case robustness levels exactly as required. The normalised efficiency decreases slightly as the robustness requirement increases due to the added costs from additional warehouses and secondary links. Complexity indices are also listed in the last column. As expected, complexity increases as robustness is required.

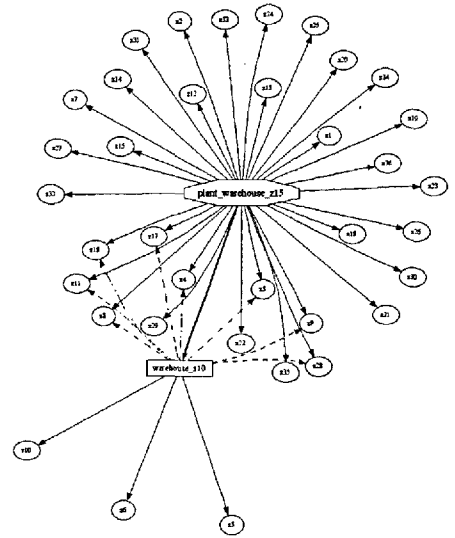
Robustness Requirement	Topology	Cost (in £)	Normalised Efficiency	Robustness		Complexity
				Worst case	Average	
$\beta_{\text{worst-case}} = 0.0$	Star	950204	1.000	0.00	0.946	1.000
$\beta_{\text{worst-case}} = 0.1$	2 Hubs	954142	0.996	0.10	0.949	1.083
$\beta_{\text{worst-case}} = 0.3$	2 Hubs	956858	0.993	0.30	0.952	1.250
$\beta_{\text{worst-case}} = 0.5$	2 Hubs	960392	0.989	0.50	0.956	1.361
$\beta_{\text{worst-case}} = 0.7$	3 Hubs	965118	0.985	0.70	0.959	1.500
$\beta_{\text{worst-case}} = 0.9$	3 Hubs	970405	0.979	0.90	0.968	1.778

Table 3.1: Logistical network performance measurement table

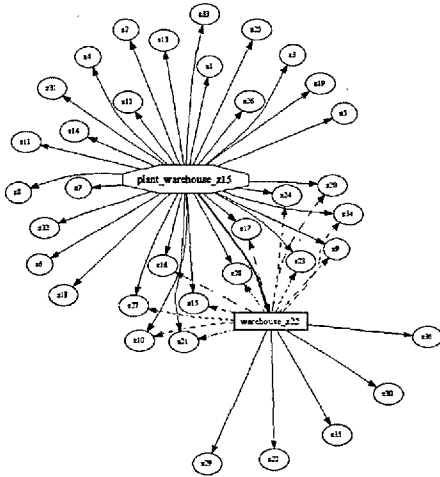
When the robustness requirement is introduced, the optimal network topology is changed. Table 3.1 shows that the number of warehouses or hubs increases with an increasing minimum robustness requirement. When there is no robustness requirement, there is only one hub, a star topology. However, for minimum robustness requirements of 0.1, 0.3, and 0.5, two warehouses or hubs are needed. For minimum robustness requirements of 0.7 and 0.9, the networks require three hubs. We present illustrations of logistical networks for each minimum robustness requirement in Figure 3.4. Primary links are displayed as solid lines, whereas secondary links are displayed as dashed lines.



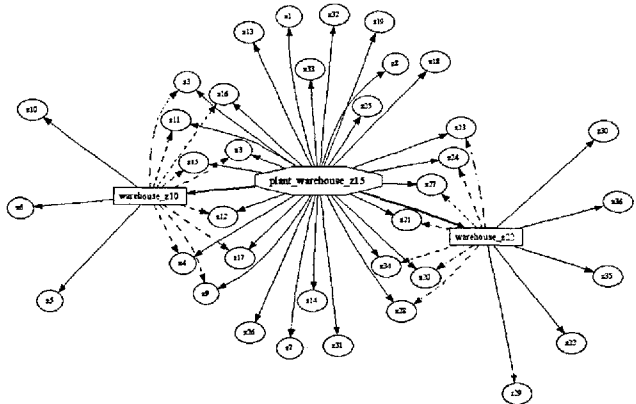
(minimum robustness = 0.1)



(minimum robustness = 0.3)



(minimum robustness = 0.5)



(minimum robustness = 0.7)



(minimum robustness = 0.9)

Figure 3.4: Logistical network structures for various robustness requirements

In all network structures, one warehouse will be at zone 15 which is the same location as the plant because it is the most cost efficient in terms of transportation cost. Additional warehouses are added to meet robustness requirement. In all cases, the additional warehouses are located in zone 10 or zone 22, or both. When selecting locations for warehouses, key factors are robustness requirement, demand profile at each location, and distance between customer zones. Therefore, depending on the combination of these factors, a warehouse location that is optimal for one configuration may not be optimal at the other configurations even though they have the same number of warehouse. For example, both robustness = 0.1 and robustness = 0.3 cases have 2 warehouses, however, the former has warehouses at zones 15 and 22 but the latter has warehouses at zones 15 and 10.

Additionally, interesting relationships that we can explore further are the relationships between efficiency, robustness, and complexity. Below we illustrate their relationships based on results from our case study in Figure 3.5.

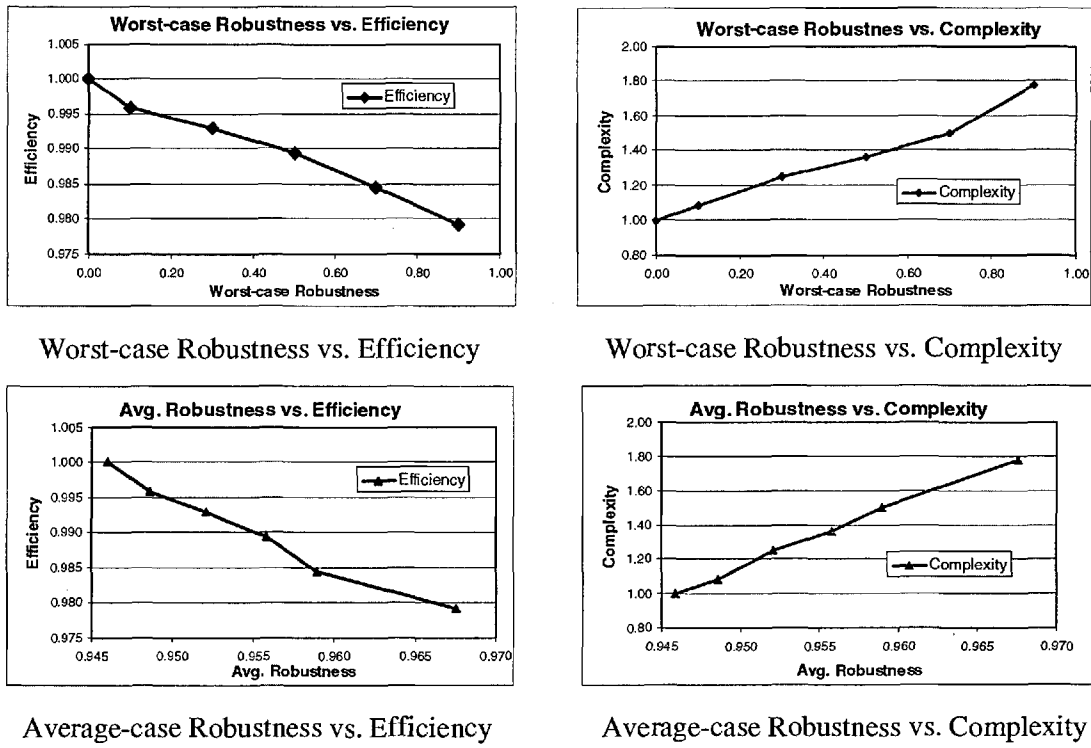


Figure 3.5: Relationships of Robustness versus Efficiency and Complexity

From the graphs above, it is clear that when robustness increases, the efficiency decreases whilst complexity increases. The trends are the same for both worst-case robustness and average-case robustness. Topological robustness is said to be achieved by the way the networks are connected

without introducing more links or components, in other words, topological robustness does not increase complexity. For example, in Figures 3.6(a) and 3.6(b) both networks have the same number of warehouses (represented by boxes), customers (represented by circles), and links (represented by lines). Assuming that all customers have equal demand, the worst-case robustness of 3.6(b) is higher than 3.6(a) because in the worst case 3.6(b) can fulfill half of total demand whilst 3.6(a) will be able to serve only one-sixth. Robustness gained in 3.6(b) is, hence, topological robustness. However, there is a limitation in topological robustness. Notice that we cannot increase the robustness of 3.6(b) further by only changing the way the network is connected without adding more links or components. Figures 3.6(c) and 3.6(d) demonstrate how the robustness of the network can be improved further to two-thirds of total demand (dashed lines represent secondary links). However, as either more links or components are added, the robustness gained is called “redundancy robustness” and, as a result, the complexity of the network increases. The graphs in Figure 3.5 show that the complexity increases as the minimum robustness requirement of the logistical network increases. Therefore, in our case, robustness is gained at least partly from redundancy, not entirely from topological structure.

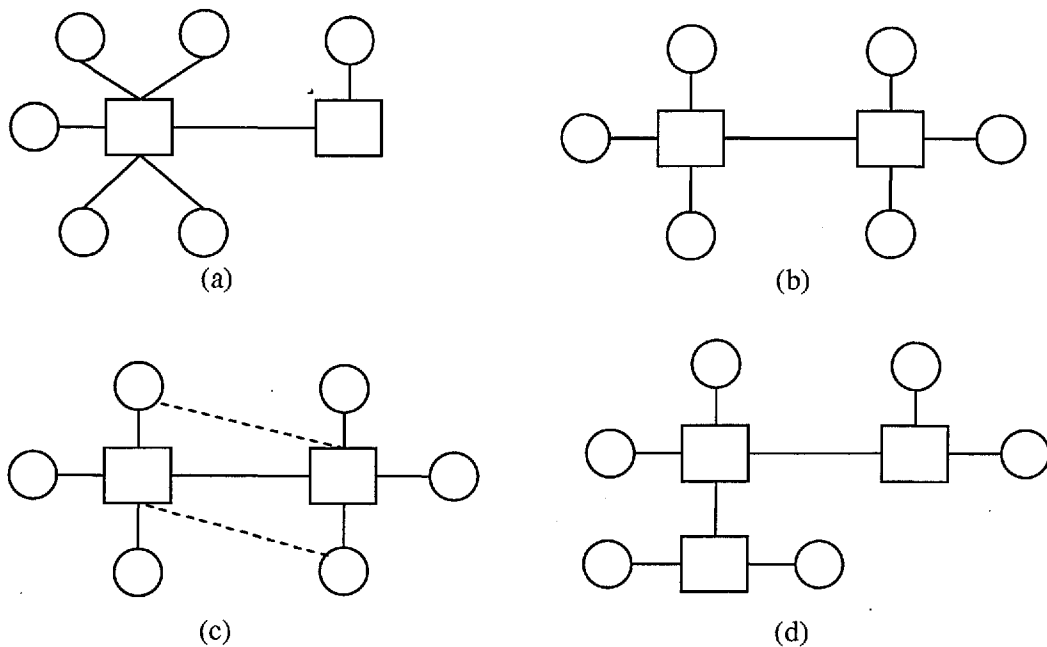


Figure 3.6: Different logistical network structures with topological or redundancy robustness

3.4 Impact of Transportation Economy of Scale

Another real-world characteristic of transportation cost structures is the economy of scale. The cost of transportation per unit product is usually lower when large volumes are shipped. We will explore how this affects the network structure and other properties.

Modified Transportation Costs

Suppose that the transportation cost rate has the following scheme. With economy of scale, the unit transportation cost rate (/ton/km) for shipment of (0, 2500] ton is £2.4, (2500, 5000] ton is £2.0, (5000, 7500] ton is £1.6, and above 7500 ton is £1.2, respectively as in the graph below.

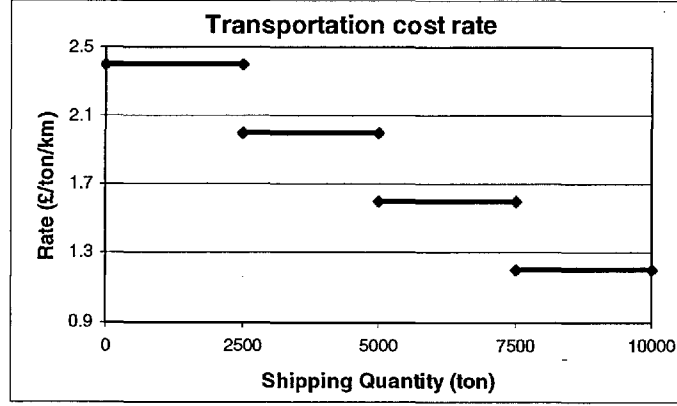


Figure 3.7: Unit transportation cost rate vs. Shipping Quantity

Now we introduce another binary variable ZP_{il} where the index i is warehouse location and index l is transportation cost rate level.

$$ZP_{il} = \begin{cases} 1 & \text{if the volume shipment from a plant to a warehouse falls into cost rate level } l \\ 0 & \text{otherwise} \end{cases}$$

where $l = 1, 2, 3, 4$.

The variable transportation cost rate C_l^i is then becomes, £2.4, £2.0, £1.6, and £1.2 when $l = 1, 2, 3, \text{ and } 4$, respectively. This translates into the following new constraints:

$$\sum_{l=1}^4 ZP_{il} = W_i, \quad \forall i \quad (3.16)$$

$$P_{il} \leq M \cdot ZP_{il}, \quad \forall i, l \quad (3.17)$$

P_{il} is the quantity of product shipped from the manufacturing plant to warehouse i , if it exists.

As a consequence, constraint (3.7) has to be replaced by constraint (3.18).

$$\sum_{j=1}^n Q_{ij} = \sum_{l=1}^4 P_{il}, \quad \forall i \quad (3.18)$$

The total transportation cost from the plant to warehouses is then:

$$\sum_{i=1}^n \sum_{l=1}^4 C_l^i \cdot P_{il} \cdot LP_i$$

There is no volume discount for shipment from the warehouse to any customer zone because the demand of each customer zone does not exceed 1,000 ton, thus the transportation cost rate from any warehouse to any customer zone is fixed at £2.4/ton/km. The objective function is then:

$$\min \sum_{i=1}^n C^w \cdot W_i + \sum_{i=1}^n \sum_{j=1}^n C^h \cdot Q_{ij} + \sum_{i=1}^n \sum_{l=1}^4 C_l^t \cdot P_{il} \cdot LP_i + \sum_{i=1}^n \sum_{i=1}^n C_l^t \cdot Q_{ij} \cdot L_{ij} + \alpha \sum_{i=1}^n \sum_{j=1}^n D_j \cdot L_{ij} \cdot XS_{ij} \quad (3.19)$$

3.5 Result of a model with economy of scale in transportation

Other details are kept unchanged from the original model, and again we can solve the above model as an MILP problem without robustness requirement (i.e. $\beta = 0$). The optimal supply chain network with economy of scale in transportation is illustrated in Figure 3.8.

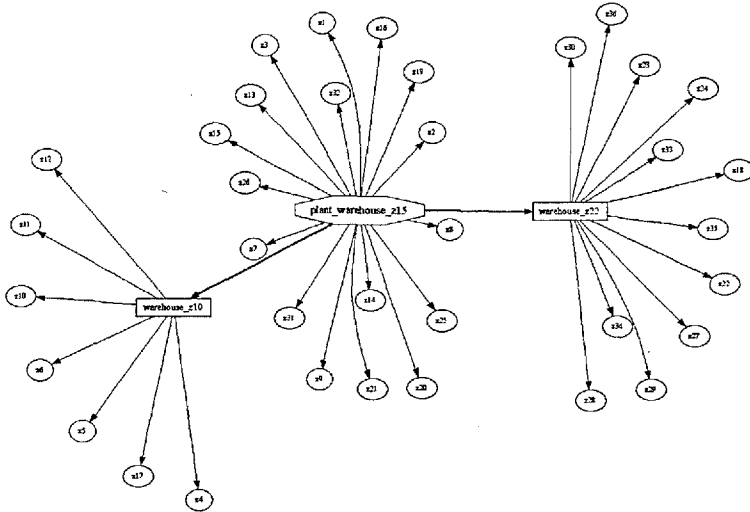


Figure 3.8: Logistical network of model with transportation economy of scale

Table 3.2 demonstrates various performance measures of the network in Figure 3.8.

Case	Topology	Cost (in £)	Normalised Efficiency	Robustness		Complexity
				Worst case	Average	
Economy of Scale	3 Hubs	1003803	0.947	0.60	0.949	1.056

Table 3.2: Performance measurement of network with scale economy

The result shows that the economy of scale impacts the optimal network topology, efficiency, robustness, and complexity. The normalised efficiency is affected by the transportation cost rate structure. In this example, the normalised efficiency is 0.947 and could be increased if transportation cost rate structure is revised (e.g. through negotiation). The worst-case robustness

is 0.60 with complexity level 1.056. Compared with the no economy of scale case where complexity level is as high as 1.361 when worst-case robustness is only 0.50, the economy of scale case is potentially more desirable. It suggests that we might be able to obtain a desirable level of efficiency and robustness without much increase in complexity by carefully designed networks employing economies of scale, possibly through long-term supply contracts.

The model with economy of scale in transportation has 32,619 constraints, 11,816 continuous variables and 10,584 integer variables, whilst the model without economy of scale has 14,439 constraints, 3,932 continuous variables and 5,256 integer variables, respectively. We programme the problem by using the General Algebraic Modelling System (GAMS) as the modelling language and CPLEX as a solver. Since modern solvers like CPLEX can handle millions of constraints and variables, our models here are well within manageable level. Running on Imperial College Linux clusters, each of our models can be solved within a few minutes.

3.6 Network topological properties

Another interesting characteristic of network topology is that, for many real-world networks such as the Internet, protein interaction networks, etc., the plot of degree of vertex (k_i) versus its probability $p(k_i)$ is often found to be a power law distribution, i.e. $p(k_i)$ is proportional to $k^{-\gamma}$ (Albert and Barabási, 2002). We produce such plot k_i against $p(k_i)$ for the above supply chain networks in Figure 3.9. Of course, for some networks there are very few points.

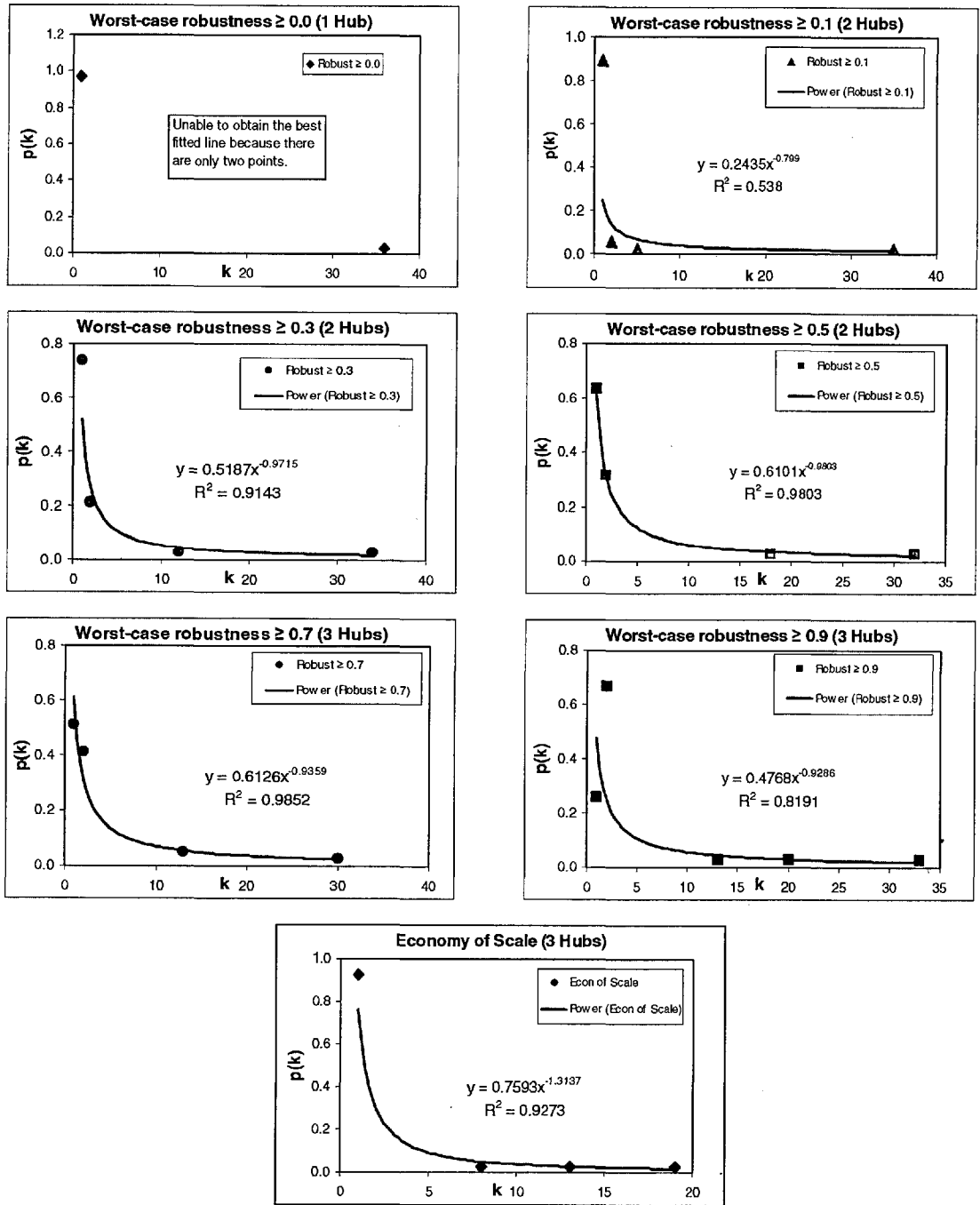


Figure 3.9: Degree distributions of logistical networks

Along with degree distribution, we calculate average path length (APL) and clustering coefficient (C) of the logistical networks. A summary of the exponent ($-\gamma$), goodness of fit (R -squared) of the best fit line, average path length, and clustering coefficient is shown in Table 3.3. For a comparison we also include average path length ($APL_{rand} = \frac{\ln(N)}{\ln(\langle k \rangle)}$, N = number of vertices) and clustering coefficient ($C_{rand} = \frac{\langle k \rangle}{N-1}$) of random networks of the same size and average degree.

Robustness Requirement	Gamma($-\gamma$)	R-squared	$\langle k \rangle$	APL	APL _{rand}	C	C _{rand}
$\beta_{\text{worst-case}} = 0.0$	N/A	N/A	1.946	3.892	5.424	0.000	0.054
$\beta_{\text{worst-case}} = 0.1$	-0.799	53.80%	2.053	4.071	5.058	0.070	0.055
$\beta_{\text{worst-case}} = 0.3$	-0.972	91.43%	2.368	4.085	4.219	0.214	0.064
$\beta_{\text{worst-case}} = 0.5$	-0.980	98.03%	2.579	4.131	3.840	0.318	0.070
$\beta_{\text{worst-case}} = 0.7$	-0.936	98.52%	2.769	4.394	3.597	0.416	0.073
$\beta_{\text{worst-case}} = 0.9$	-0.929	81.91%	3.282	4.108	3.083	0.674	0.086
Econ of Scale	-1.314	92.73%	1.949	5.274	5.491	0.000	0.051

Table 3.3: Summary of topological properties of logistical networks

Even though the numbers of data points are quite small, these logistical networks closely follow the power-law distribution except the second network (worst-case robustness ≥ 0.1). Networks whose average path lengths are small are classified as *small-world networks*. In a small-world network, most vertices can be reached from every other by traversing along a small number of edges. Many natural and man-made networks have been characterised as small-world networks. The logistical networks also exhibit the “small-world property” as their average path lengths are quite small compared to the number of vertices. The average path length is smallest when there is no robustness requirement. It is also generally known that the average path length of a random network is small as it scales as the logarithm of the number of vertices. However, excluding the $\beta_{\text{worst-case}} = 0$ and economy of scale cases for which no secondary links are present, the clustering coefficients of the logistical networks are all larger than those of random networks. Many other real world networks also exhibit greater clustering coefficients than random networks of the same size (Albert and Barabási, 2002). Notice that the more secondary links, the higher the minimum robustness, and the larger the clustering coefficient. This probably suggests that a high clustering coefficient positively correlates with the minimum robustness. It is also well-known that random networks often leave some of their vertices disconnected, and hence have low clustering coefficients and presumably low robustness. However, our results also show that it is not always the case that robustness is low when the clustering coefficient is low as is evident in the transportation economy of scale case. Therefore, we might assume a high clustering coefficient as only a potentially sufficient but not necessary condition for logistical network robustness.

Complex networks that follow the power-law degree distribution are often called “scale-free networks”. They are called so because their degree distributions are functions of degree of vertex (k) and its exponent ($-\gamma$) and are entirely independent of scale or number of vertices in the networks. The negative exponents ($-\gamma$) suggest that the probability of degree of vertex decays at a power rate of γ as degree of vertex increases. This feature is a characteristic of hub-like

network structures. Hence, from the results shown in Figure 3.9 and Table 3.3, supply chain managers can expect robust logistical networks to be hub-like with only a few highly connected components. This can help in giving the managers a rough idea of what their logistical networks should look like and estimating number of links required. The comparatively low average path length implies that flow of product or information along linkages in robust logistical networks from one component to any other component on average requires only few steps of transferral, though the average path length increases with the robustness requirement.

3.7 Managerial Implications

To protect against undesirable disruption and to gain long-term survival, logistical network robustness should be one of the important characteristics that supply chain managers consider when they design a supply chain. To achieve a desired level of robustness, they can construct a mathematical model according to their constraints and objectives including robustness requirements such as the one we show in our case study. The mathematical model can answer questions such as how many warehouses should be established, which warehouse should be assigned to particular customers, where should secondary links be established etc. However, there are other managerial implications as well. For example, when a manager establishes secondary links in his or her logistical network he or she has to make sure that the warehouses have enough excess capacity to accommodate secondary links. If this will incur cost, it can be easily incorporated into a mathematical model.

The normalised efficiency should also be closely analysed. An efficiency index of 0.5 means a company has to pay twice what it might otherwise pay to operate the most efficient logistical network. A supply chain manager should check the reasonableness of the efficiency of the logistical network derived by the mathematical model. In our framework, logistical network complexity is not directly incorporated into the mathematical model, however, it is measured after the network topology has been realised. A manager should therefore examine whether the network complexity is reasonable for his or her supply chain. The complexity index of 2 means a supply chain manager has to maintain twice as many logistical channels (though not all may be active) of what would otherwise be in the most efficient logistical network. If this is not reasonable, the manager should review whether he or she requires too high a level of robustness or sets too low a cost of a secondary link. This iterative process of logistical network design would yield the most desirable, both for short term and long term, logistical network.

Particular supply contracts may help to increase logistical robustness without much decrease in efficiency or increase in complexity. Supply chain managers should be aware of this potential benefit and take this into account when negotiating supply contracts.

3.8 Summary and Discussion

In this chapter we define logistical network robustness and argue that supply chain managers should consider logistical network robustness when they design a logistical network. We also present measures for logistical network complexity and normalised efficiency. Our proposed concept is demonstrated through a mathematical model and a case study to design a logistical network formulated using Mixed-Integer Linear Programming. The objective function is to minimise cost (absolute efficiency) subject to a minimum robustness requirement. After solving the mathematical model, we obtain different logistical network configurations which satisfy different robustness requirements with the lowest logistical network operating cost. The normalised efficiency and complexity between network configurations are then compared. From our results, when the robustness requirement increases, logistical network complexity also increases whilst the network efficiency decreases. Supply chain managers have to trade-off robustness with both complexity and efficiency. There is no absolute rule what the robustness requirement should be. Supply chain managers need to use their judgement as to what is a desirable robustness level, possibly from the worst-case service level requirement. We also make a small modification in the model to accommodate economy of scale in the transportation cost rate. The result shows that we might be able to obtain a desirable level of robustness without much increase in complexity or decrease in efficiency through carefully designed networks, possibly inducing economies of scale and long-term supply contracts. Finally, we show that robust logistical networks in this chapter typically follow the power-law degree distribution, a well-known topological property of scale-free networks, and exhibit small average path lengths. When secondary links are present, the logistical networks also have large clustering coefficients.

Although the mathematical model in our case is fairly simple, it demonstrates our concept well and reveals the relationship between the three performance measurements. The model can also easily be modified to accommodate other constraints and requirements for more complex supply chain environments.

Chapter 4

Vehicle Routing Networks and Complexity

Distribution (vehicle) scheduling is one of the most fundamental decisions in supply chain management. The previous chapter discusses the design of a logistical network. In this chapter we discuss planning of vehicle routing and scheduling in different logistical networks developed in chapter 3, measurement of a vehicle routing network's complexity over various time horizons, and consistency of vehicle routing network structure to the power law and some statistical measures of network properties.

4.1 Introduction

Two important operational details in supply chain management are production scheduling and distribution scheduling. Distribution scheduling is often achieved by vehicle routing and scheduling. Vehicle routing and scheduling assigns the customers to be visited by a particular vehicle and the sequence in which they will be visited. Logistics design in a supply chain usually starts from the design of a logistical network. Once the logistical network has been fixed, it establishes constraints and becomes an environment within which vehicle routing and scheduling will be determined. The objective of vehicle routing decision making is usually to minimise distribution costs. One way to plan vehicle routing is using an optimisation technique to obtain the optimal solution. Optimising vehicle routing is often related to the Travelling Salesman Problem (TSP). It is well known that when an exact optimisation algorithm is used to find the optimal solution, the difficulty of travelling salesman problems increases rapidly (e.g. the time required to solve the problem is prohibitively long) as the number of locations to be visited increases (Hillier and Lieberman, 2005). A more popular way of planning vehicle routes is using heuristic methods which generally provide a good solution within a reasonable time.

There are several papers which propose heuristic algorithms for vehicle routing (for example Lee and Ueng, 1999; Tan *et al.*, 2001; Ropke and Pisinger, 2006). Generally, the objective is to minimise cost, time, distance etc. Research in vehicle routing has been conducted by researchers in a wide range of academic disciplines. This is because the vehicle routing problem is highly relevant to the real world of logistics management (Sateesh and Ray, 1992; Eibl *et al.*, 1994). The majority of vehicle routing studies have focused on identifying appropriate vehicle routing algorithms, with the most popular techniques being Tabu Search (Montané and Galvão, 2006; Wassan, 2006), Simulated Annealing (Czech and Czarnas, 2002; Li and Lim, 2003), and Genetic Algorithms (Hwang, 2002; Berger and Barkaoui, 2004; Pankratz, 2005a, 2005b), or a combination of these techniques or combinations with exact algorithms.

However, from an operations point of view, vehicle routing and scheduling is not only about finding optimal solutions but also about executing vehicle routing according to the schedule as well. The difficulty of executing vehicle routing should be taken into account when a supply chain manager plans vehicle routings, selects vehicle types, and, if possible, designs logistical networks. Vehicle routings can become very complex no matter which method or algorithm is used for planning. Complexity in vehicle routing is usually referred to as computational complexity or solvability of the vehicle routing problem such as whether polynomially solvable or NP-hard (for example see de Paepe *et al.*, 2004; Archetti *et al.*, 2005; Hassin and Rubinstein, 2005). However, complexity from an *operations* point of view is rarely described in the literature.

The entropic approach to complexity measurement has been applied to many kinds of systems. Yu and Efstathiou (2002) proposed a complexity measurement procedure for a production line network using the entropic approach. Their procedure successfully solves the state-dependent limitation of the conventional method. Wu *et al.* (2002) presented a simulation study on the operational complexity of inventory and delivery quantity in the manufacturing industry. Entropic complexity defines complexity in terms of the expected amount of information required to describe the state of the system or entropy in the system, which can be interpreted as the level of communication and coordination taking place in the system. There are a number of research papers on complexity measurement based on the entropic approach in production aspects of supply chain systems (for example Calinescu *et al.*, 1997). We argue that a similar approach may be applicable to vehicle routing since we can define different delivery routes as different states of the system.

In vehicle routing, when each location is linked according to vehicle routes, we can also consider this as a graph or a network. Natural or man-made networks often hold certain distinct

properties such as the power-law degree distribution (Albert and Barabási, 2002). Therefore, treating them as one type of network, we can also investigate the topological properties of vehicle routing networks as well. In this chapter we discuss planning of vehicle routing and scheduling in different logistics configurations with two different demand profiles. We perform measurements of vehicle routing entropic complexity over various time horizons, and examine the consistency of vehicle routing network topology to the power-law degree distribution and statistical measures of network topological properties (i.e. average and variance of degree of vertex).

4.2 Problem description

In our investigation into vehicle routing complexity and vehicle routing network topological properties, we consider the logistical networks presented in chapter 3 where customers are located in the following customer zones denoted by number 1 to 36.

Customer Zone

1	2	3	4	5	6
7	8	9	10	11	12
13	14	Plant	16	17	18
19	20	21	22	23	24
25	26	27	28	29	30
31	32	33	34	35	36

Figure 4.1: Customer zone geography

There is one plant producing products for all customers which is located at zone 15, as shown in Figure 4.1. The size of each zone is 10 km vertically and 10 km horizontally. The distance between zones is assumed to be the direct Euclidean distance between centres of each zone. From the customer zones as shown in Figure 4.1 and customers' demand profile, in chapter 3, we propose a logistical network design procedure with complexity and robustness considerations to determine how many warehouses should be built and where they should be located. We obtain different logistics configurations for different robustness requirements and whether there is economy of scale in transportation cost as summarised in Table 4.1.

Configuration	Economy of Scale?	Robustness requirement	Number of Warehouse	Location of Warehouse	Customers Served
A	No	None	1	zone 15	All customer zones
B	No	10%	2	zone 15 zone 22	All except zone 29 and 36 zone 29 and 36
C	No	30%	2	zone 10 zone 15	zone 5, 6, 10 All except zone 5, 6, 10
D	No	50%	2	zone 15 zone 22	All except zone 22, 29, 30, 35, 36 zone 22, 29, 30, 35, 36
E	No	70%	3	zone 10 zone 15 zone 22	zone 5, 6, 10 The rest of customer zones zone 22, 29, 30, 35, 36
F	No	90%	3	zone 10 zone 15 zone 22	zone 5, 6 The rest of customer zones zone 22, 29, 30, 35, 36
G	Yes	None	3	zone 10 zone 15 zone 22	zone 4-6, 10-12, 17 The rest of customer zones zone 18, 22-24, 27-30, 33-36

Table 4.1: Logistics Configuration

The customers served by each warehouse as shown in Table 4.1 are based on primary links between warehouses and customers which will be the warehouse-customer assignments that we will consider for vehicle routing.

4.3 Vehicle routing procedure

Obtaining exact optimal solutions for the vehicle routing problem generally requires a large computational time due to its combinatorial nature; the number of feasible solutions increases exponentially with the number of customers to be visited. Therefore, a heuristic method which is a procedure that is likely to find a good feasible solution, but not necessarily an optimal solution, is often used in the vehicle routing problem. There are various heuristic algorithms suggested in the literature. For our study we adapt a simple but proven to provide relatively good result heuristic method called “Generation by Node insertion” method presented by Jalisi (2001) as described in the procedure below:

For each warehouse

Step 0: A vehicle starts its journey from a warehouse. The vehicle can only travel to the customers assigned to the warehouse by the logistics configuration as per Table 4.1.

Step 1: A journey starts as a vehicle travels to the farthest customer whose demand has not been filled. If there is a tie in the farthest customer, select any of the ties.

If Demand of the first customer > vehicle capacity
Then Unload product equal to vehicle capacity
 Reduce demand of the first customer by vehicle capacity
 Go back to the warehouse and go to *Step 0*
Else Reduce current vehicle capacity by demand of the first customer
 Available time = Maximum Travel Time – (Travel Distance / Vehicle Speed) –
 Unloading Time
 Goto *Step 2*

Step 2: If capacity and time is allowed, a vehicle travels on to fill the demand of the nearest unfulfilled customer from its current location. If there is a tie in the nearest unfulfilled customer, select any of the ties.

If Available time > 0 and current vehicle capacity > 0
Then Find the nearest customer whose demand is unfilled and $\leq$ current vehicle capacity
 If All customer demands are filled
 Then Go back to the warehouse and *Stop*
 ElseIf All customer whose demand is unfilled has demand > current vehicle capacity
 Then Go back to the warehouse and start *Step 0*
 ElseIf Travelling to the nearest customer whose demand is unfilled and $\leq$ current
 vehicle capacity requires more than available time
 Then Go back to the warehouse and start *Step 0*
 Else Reduce current vehicle capacity by demand of this customer
 Available time (hr) = current available time – (Travel Distance / Vehicle Speed)
 – Unloading Time
 Repeat *Step 2*
Else Go back to the warehouse and start *Step 0*

In addition, we assume that there will always be enough vehicles for vehicle routing. Therefore, the number of available vehicles is not a constraint in our problem.

4.4 Complexity measure

The complexity measured in our study is the structural complexity based on an entropic or information-theoretic complexity measure (Calinescu *et al.*, 2000), which is a complexity of planned activity. The entropic complexity is a relevant choice of measure in this case because it

reflects the expected amount of information required to describe the state of the vehicle routing operations which implies the amount of coordination and planning effort. The entropic complexity measure based on an information-theoretic point of view is defined as follows;

$$H_E = -\sum_{i=1}^O \sum_{j=1}^D p_{ij} \log_2 p_{ij} \quad (4.1)$$

where H_E is the quantity of entropic complexity

O is total possible number of origin points of a journey

D is total possible number of destinations of a journey

p_{ij} is the probability of a journey being on route from location i to location j

In this case, the state of the system is the state of a vehicle being on a journey from zone i to zone j . The probability of each state is then defined as the fraction of time out of total travel time in the logistics network that a vehicle is in that state. From the entropic complexity calculation formula, the complexity will be low when most of the time the vehicles follow a small number of routes. The complexity will be high if the vehicles take many different routes with equally likely probability. For example, suppose there are 10 possible origins and 10 possible destinations, hence, there are 100 possible routes. If it is equally likely that the vehicle will be in any particular route, all p_{ij} will equal $\frac{1}{100}$ or 0.01. Therefore, the entropic complexity of this vehicle routing system is

$$-\sum_{i=1}^{10} \sum_{j=1}^{10} 0.01 \times \log_2 0.01 \approx 6.644$$

4.5 Degree distribution

In vehicle routing, the vehicle travels between each zone within the customer zones to deliver the product. If we look at vehicle routing in terms of a graph, we can view each zone as a vertex and a route that the vehicle travels between each zone as an edge between the two vertices. There will be an edge between any pair of vertices only when the route between them is included in vehicle routing, otherwise there is no edge between them. We will call the resulting graph obtained from the solution to the vehicle routing problem a vehicle routing network. According to a graph theoretic framework, vertex degree or degree (k) of a vertex is the number of edges attached to it, $n(k)$ is the number of vertices with degree k in the graph, and $p(k)$ is the probability that any vertex will have degree k (Deo, 1974). It is widely known that, for many real-world networks, the plot of degree of vertex (k) versus its probability $p(k)$ is often found to

be a power distribution, i.e. $p(k)$ is proportional to $k^{-\gamma}$ (Albert and Barabási, 2002). The power-law degree distribution can be observed not only in physical networks such as the World Wide Web (Albert *et al.*, 1999), metabolic networks (Jeong *et al.*, 2000), etc. but also in ecological relationships such as food webs (Montoya and Solé, 2000) and in personal relationship such as the movie actor collaboration network (Albert and Barabási, 2000) and science collaboration networks (Barabási *et al.*, 2001), etc. In this chapter, we will examine the degree distribution of vehicle routing network as well as the statistical measures (mean and variance) of the degree of vertex.

4.6 Computational experiments: A case study

To apply our approach and to perform computational experiments, we consider the case where the customer demands are stochastic, so this belongs to the class of vehicle routing problems with stochastic demands (Chepuri and Homem-De-Mello, 2005). Based on the vehicle routing procedure for daily delivery mentioned above, computational experiments are performed for each logistics configuration in Table 4.1. The simulation and analysis program for this study is coded in *Visual Basic* because it can handle spreadsheet data very well. Additional settings to complete the experiment are detailed below.

Daily Demand Profiles

There are two different daily demand profiles in this case study. The demands of a given customer zone during each period are modelled as independent and identically distributed random variables. The first profile is when daily demand at each zone follows a uniform distribution ranging from 0 to $2 * \text{Avg.Monthly.Demand} / 30$ units where the average monthly demand of each zone is displayed in Figure 4.2.

393.5	45.5	744.9	813.6	919.6	559.0
158.6	153.4	618.8	882.3	78.7	706.5
973.4	452.5	358.1	440.6	593.0	321.3
624.4	301.8	321.0	821.0	913.8	399.2
542.1	374.9	447.0	30.4	751.6	701.8
202.4	793.3	829.3	942.0	899.9	303.1

Figure 4.2: Average monthly demand profile for each customer zones

The second profile is when daily demand at each customer zone is generated randomly based on a uniform distribution ranging from 0 to $2 * 1000 / 30$ units (demand distributions of all customer zones are homogeneous).

Vehicle Capacity

Three types of vehicle are considered in this computational experiment. They are different in maximum load capacity: 240 units (large truck), 120 units (medium truck), and 60 units (small truck), respectively.

Maximum Travel Time

The maximum travel time for each truck is set at 7 hours. A vehicle must arrive at the last destination of each journey within 7 hours after leaving a warehouse.

Vehicle travelling speed

Travelling speed is assumed to be constant at 40 km/hr. The travelling time to each destination is then calculated as $\frac{\text{Travelling Distance}}{40}$.

Unloading Time

At each customer a vehicle travels to, 30 minutes (0.5 hour) of unloading time is assumed regardless of the drop size. This unloading time is added to the vehicle travel time. The unloading time will not be included when we calculate state probabilities. However, it reflects the real situation as the vehicles that visit more customers will have lower effective travel times than the vehicles that visit fewer customers.

Planning Horizon

The planning horizon can be arbitrarily selected. In this case study, we experiment with planning horizons of 1, 5, 15, 50, 100, 200, 300, 400, 500, 600, 800, 1000, 1500, 2000, 3000, 4000, and 6000 days, respectively. Though planning horizons of more than a month look unrealistic, we do so to observe the steady state complexity in our case.

4.7 Simulation results

Examples of vehicle routing results for a typical day in logistics configuration A are illustrated in Figure 4.3.

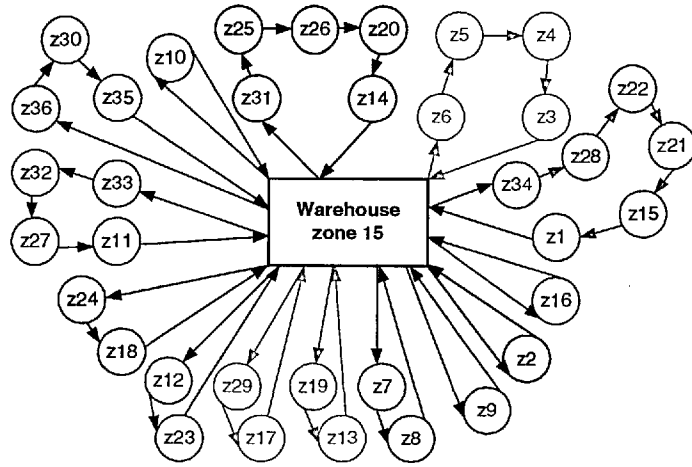


Figure 4.3: Example of vehicle routing (centralised distribution)

Typically, vehicle routing and scheduling decisions are made within short planning horizons (e.g. for the next few days), and usually after customer demand has been confirmed. We first examine the case when vehicle routing and scheduling is done for 5 days (e.g. agreement with customers to confirm demand 5 days in advance and/or agreement with a truck company to submit delivery plan 5 days in advance). Computational experiments were performed for each configuration in Table 4.1.

4.7.1 Results for First Demand Profile

900 computational experiments (with different sampled demands in each experiment) with a 5-day planning horizon are run for each configuration to obtain reliable statistics. In each simulation run, the probability that a vehicle being in any particular route is calculated using the fraction of time out of total travel time over 5-day period that a vehicle is in that route. The results are summarised in Table 4.2 below.

Configuration	Structure	Complexity Level					
		Large Truck		Medium Truck		Small Truck	
		Average ($\bar{X}$)	STD (σ_x)	Average ($\bar{X}$)	STD (σ_x)	Average ($\bar{X}$)	STD (σ_x)
A	1 Hub	5.202	0.090	6.184	0.097	6.528	0.065
B	2 Hubs	5.236	0.095	6.130	0.101	6.480	0.063
C	2 Hubs	5.191	0.056	6.024	0.104	6.460	0.071
D	2 Hubs	5.130	0.037	6.100	0.087	6.437	0.059
E	3 Hubs	5.142	0.016	5.921	0.102	6.338	0.071
F	3 Hubs	5.264	0.075	6.099	0.098	6.466	0.065
G	3 Hubs	5.228	0.010	5.828	0.097	6.294	0.071

Table 4.2: Summary of complexity level (1st demand profile)

From Table 4.2, the standard deviation is much smaller than the average complexity. Therefore, for each configuration the complexity distribution is quite narrow. This implies that although the demand is generated randomly for each run, the vehicle routing complexity of each configuration is relatively consistent. The average vehicle routing complexities of each configuration are also charted as shown in Figure 4.4.

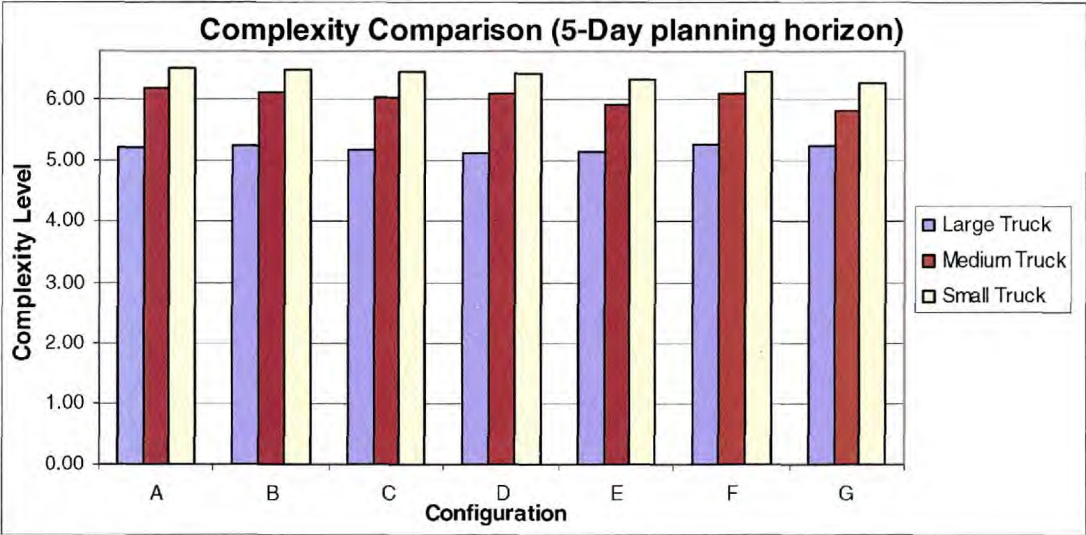


Figure 4.4: Average complexity vs. Logistics configuration (1st demand profile)

Notice that the level of complexity of vehicle routing increases as vehicle size (load capacity) decreases. This is because the smaller the load capacity, the less likely that the vehicles repeat the same routes each day. For medium and small trucks, the results show that configuration A where there is only one warehouse (centralised distribution) has the highest complexity. Centralised distribution implies that a vehicle can travel from a warehouse (the only warehouse) to any customer and from one customer to any other customer. Therefore, complexity tends to be high because there are many possible vehicle routes. Notice also that complexity is likely to decrease as the delivery becomes decentralised (multiple warehouses). However, for large trucks, the complexity level of vehicle routing at each configuration becomes comparable and there is no obvious trend. The probable explanation is that when the vehicle size is large enough (compared to average demand at each customer zone) there is a reasonable chance that vehicles repeat the same routes, and make the difference in logistics configuration less influential. We also plot complexity level of each logistics configuration and vehicle size against varying planning horizons. The results of the computational experiments are plotted and shown in Figure 4.5.

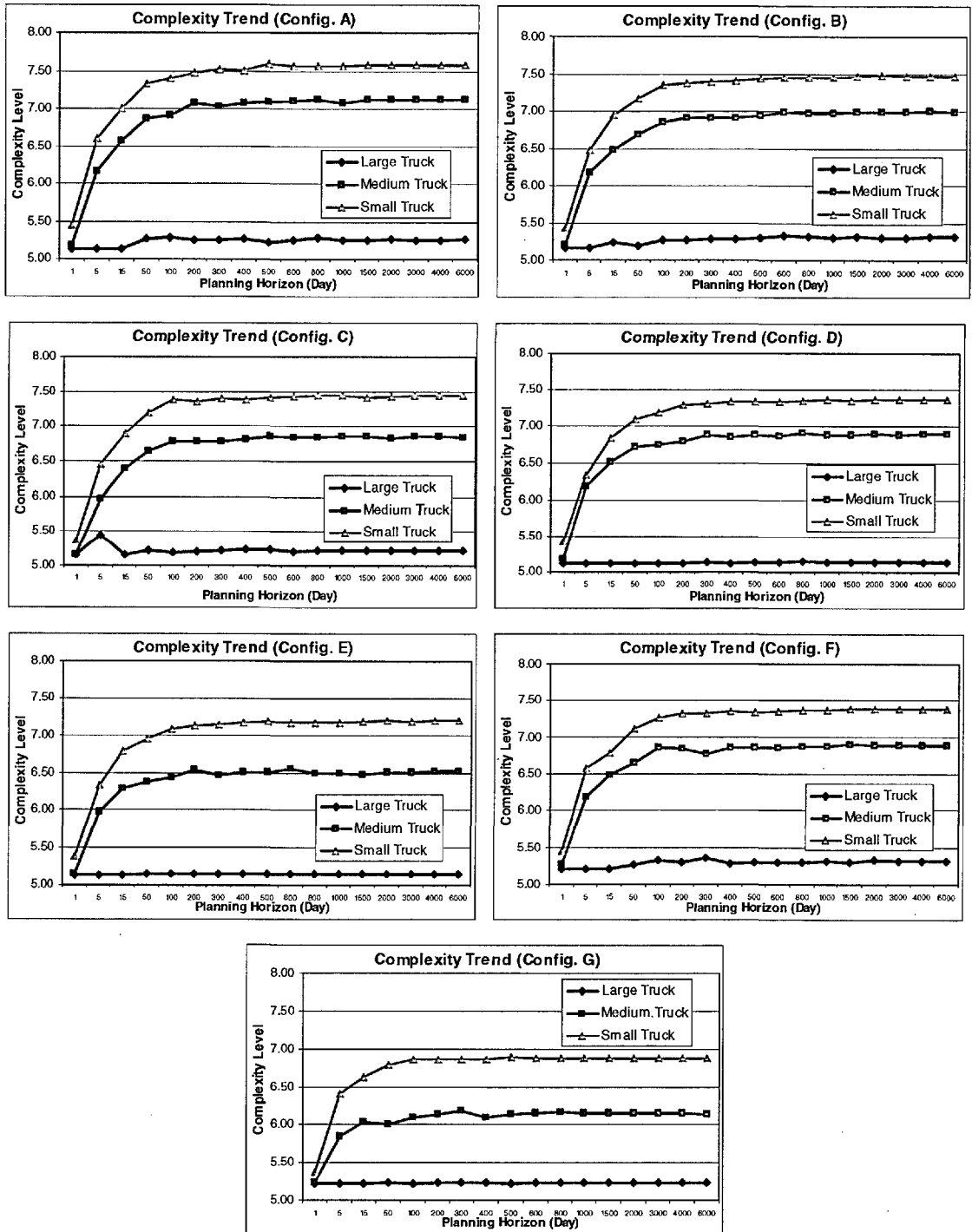


Figure 4.5: Complexity level vs. Planning horizon (1st demand profile)

There are some common patterns in the results above. Firstly, other things being equal, the smaller the vehicle the greater the complexity rule applies. Secondly, for medium and small trucks, complexity increases as the planning horizon increases and after a certain length of time (around 200 days in this case) the system reaches a steady state. For large trucks, complexity either increases slightly or remains stable as the planning horizon increases. Thirdly, for medium and small trucks, the trend that complexity level decreases as degree of centralisation

decreases (from configuration A to G) holds for all planning horizons, though this is not always true for large trucks. These trends can only be realised after computational experiments and the results cannot be generalised to other cases (e.g. different demand profiles).

For each configuration, we also record the number of trucks required for vehicle routing in each day. Figure 4.6 demonstrates an example of the plot of number of trucks required versus planning day (in configuration A with small trucks).

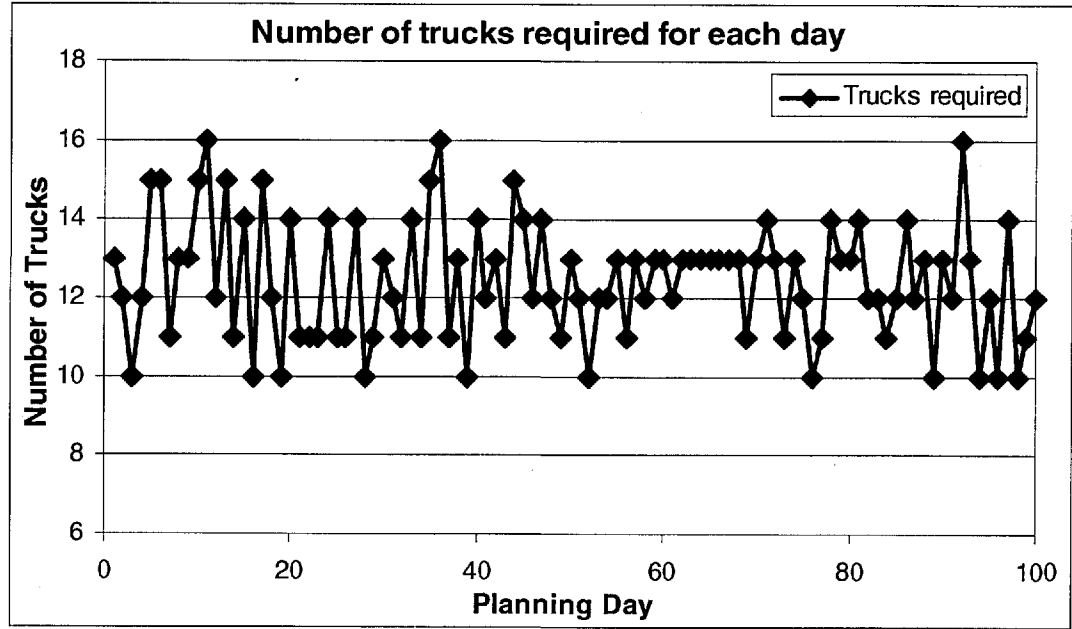


Figure 4.6: Number of trucks required for each day in configuration A with small truck

We also calculate the average, maximum, minimum, and standard deviation of the number of trucks required, and complexity level of 6000-day planning for each logistics configuration as shown in Table 4.3.

Number of trucks required and complexity level: 1st Demand Profile

Configuration	Truck Capacity	Average	Max	Min	STD	Complexity
A	Large	5.00	5	5	0.00	5.26
	Medium	6.18	9	5	0.65	7.11
	Small	12.21	17	8	1.44	7.59
B	Large	5.05	6	5	0.21	5.31
	Medium	6.86	9	5	0.62	6.99
	Small	12.65	17	8	1.43	7.47
C	Large	5.94	6	5	0.23	5.22
	Medium	6.60	10	5	0.63	6.87
	Small	12.63	18	8	1.48	7.45
D	Large	5.00	5	5	0.00	5.14
	Medium	6.61	9	5	0.75	6.90
	Small	12.75	18	8	1.47	7.37
E	Large	6.00	6	6	0.00	5.15
	Medium	7.04	10	6	0.78	6.52
	Small	13.21	19	8	1.52	7.19
F	Large	6.00	6	6	0.00	5.20
	Medium	7.22	10	6	0.75	6.60
	Small	13.19	19	8	1.52	7.24
G	Large	5.01	6	5	0.08	5.23
	Medium	7.20	10	5	0.81	6.16
	Small	13.27	19	8	1.53	6.88

Table 4.3: Statistics of number of trucks required and complexity level (1st demand profile)

From Table 4.3, the average and standard deviation of the number of trucks required and complexity level all increase as vehicle capacity decreases. Therefore, decreasing vehicle capacity does not result only in a greater number of vehicles required, but also in increasing the variability of required number of vehicles.

4.7.2 Results for Second Demand Profile

Similar to the first demand profile, we run 900 computational experiments with a 5-day planning horizon for each logistics configuration subject to the second demand profile. The results are summarised in Table 4.4.

Configuration	Structure	Complexity Level					
		Large Truck		Medium Truck		Small Truck	
		Average ($\bar{X}$)	STD (σ_x)	Average ($\bar{X}$)	STD (σ_x)	Average ($\bar{X}$)	STD (σ_x)
A	1Hub	6.087	0.116	6.510	0.078	6.590	0.044
B	2 Hubs	6.023	0.110	6.457	0.073	6.567	0.044
C	2 Hubs	6.048	0.104	6.433	0.076	6.532	0.043
D	2 Hubs	6.055	0.092	6.398	0.074	6.513	0.043
E	3 Hubs	5.959	0.089	6.326	0.075	6.455	0.046
F	3 Hubs	6.004	0.103	6.414	0.071	6.547	0.041
G	3 Hubs	5.842	0.096	6.275	0.076	6.421	0.046

Table 4.4: Summary of complexity level (2nd demand profile)

Again, the standard deviation is much lower than the average complexity. Therefore, for the second demand profile, the distribution of vehicle routing complexity of each configuration is also narrow. The average vehicle routing complexity of each configuration is charted as shown in Figure 4.7.

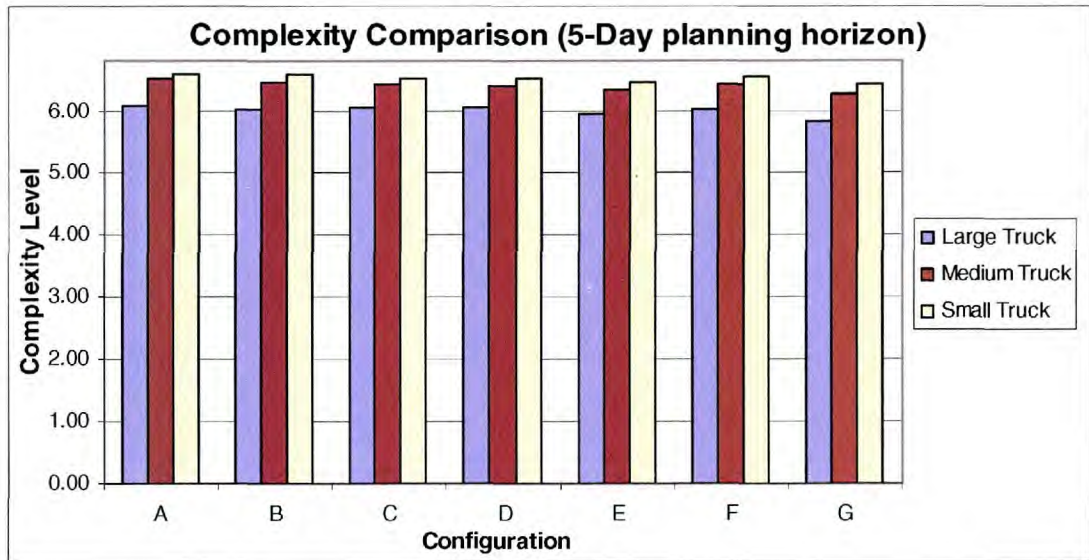


Figure 4.7: Average complexity vs. Logistics configuration (2nd demand profile)

In the second profile, the level of complexity of vehicle routing also increases as vehicle size (load capacity) decreases. In general, this result is similar to the first demand profile. However, as we will see later, this trend is only true for this planning horizon. In other cases (e.g. longer planning horizons), this trend is not always true. We also examine the complexity level of each logistics configuration and vehicle size against a varying planning horizon as for the first profile. The results of computational experiments are shown in Figure 4.8.

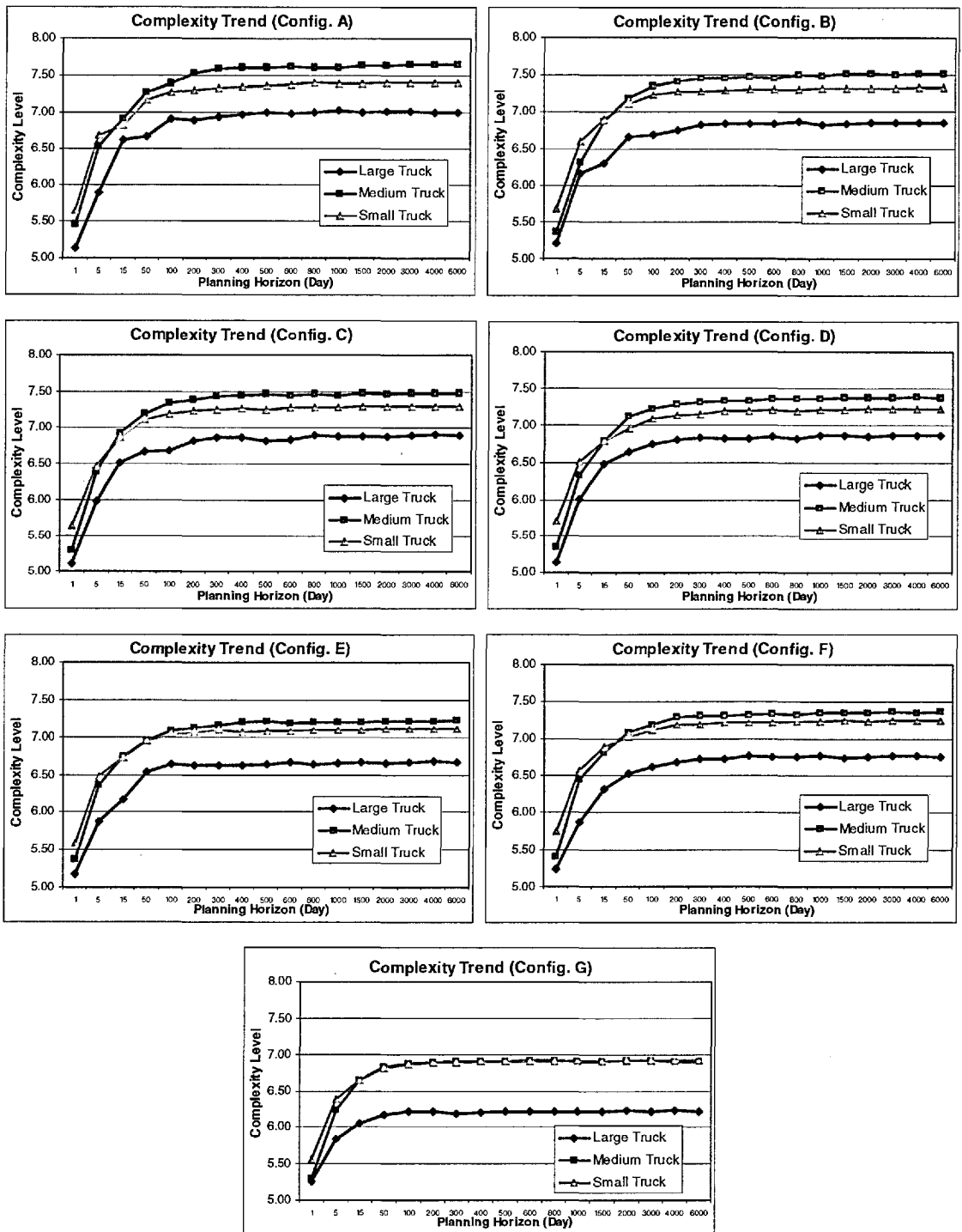


Figure 4.8: Complexity level vs. Planning horizon (2nd demand profile)

There are also some common patterns that we can see from the results above (but not all the same as the first profile). The smaller the vehicle the greater the complexity rule applies only for short planning horizons (e.g. less than 15 days). For longer planning horizons, the medium size vehicle case has a higher complexity than the small one, although the gap reduces as the degree of centralisation decreases (multiple warehouses). Again, these trends can only be realised after computational experiments and the results cannot be generalised to other cases. Furthermore,

the complexity increases as the planning horizon increases for all types of vehicle but only to some extent; after some planning horizon (around 200 days) the system reaches steady state. Finally, the trend that complexity level decreases as degree of centralisation decreases holds for all planning horizons and for all types of vehicle.

The average, maximum, minimum, and standard deviation of the number of trucks required, and complexity level of 6000-day planning for each logistics configuration subjected to the second demand profile are shown in Table 4.5.

Number of trucks required and complexity level: 2nd Demand Profile						
Configuration	Truck Capacity	Average	Max	Min	STD	Complexity
A	Large	5.71	8	5	0.55	7.02
	Medium	11.23	16	7	1.19	7.65
	Small	23.17	33	15	2.48	7.42
B	Large	6.43	8	5	0.53	6.87
	Medium	11.65	17	7	1.18	7.53
	Small	23.71	33	15	2.49	7.34
C	Large	6.29	8	5	0.48	6.90
	Medium	11.60	17	7	1.25	7.48
	Small	23.74	35	15	2.51	7.29
D	Large	6.03	9	5	0.53	6.87
	Medium	11.68	17	8	1.23	7.40
	Small	23.83	33	15	2.50	7.23
E	Large	6.57	9	6	0.57	6.68
	Medium	12.08	17	8	1.24	7.22
	Small	24.28	35	15	2.55	7.12
F	Large	6.72	9	6	0.56	6.72
	Medium	12.09	17	8	1.18	7.28
	Small	24.30	36	15	2.57	7.17
G	Large	6.58	9	5	0.70	6.23
	Medium	12.16	17	8	1.26	6.91
	Small	24.47	34	16	2.59	6.92

Table 4.5: Statistics of number of trucks required and complexity level (2nd demand profile)

From Table 4.5, the average and standard deviation of the number of trucks required increase as vehicle capacity decreases. However, contrary to the first demand profile, the complexity level does not always increase as vehicle capacity decreases.

To further investigate the relationship between the number of trucks required and the complexity level, we consider a set of extreme cases where the number of trucks is fixed. We

assume that the truck capacity is always sufficient to serve demand from the customer zones where it is assigned. Therefore, this case is planning horizon independent, demand independent, and truck-capacity independent. When the number of trucks is n , each truck will serve exactly $36/n$ customer zones. The i^{th} truck will start serving from customer zone $\frac{36(i-1)}{n} + 1$ until zone $\frac{36i}{n}$. We calculate system complexity level and average complexity level of individual truck for all possible n where 36 is divisible by n . We calculate the complexity level for each case. The average complexity level of each individual truck is calculated by dividing system complexity by the number of trucks used. The results are shown in Table 4.6.

Number of Trucks	System Complexity	Average Individual Complexity
1	4.80	4.80
2	4.88	2.44
3	4.95	1.65
4	4.98	1.24
6	5.16	0.86
9	5.21	0.58
12	5.35	0.45
18	5.55	0.31
36	6.03	0.17

Table 4.6: Complexity level in fixed route vehicle routing

The plot of complexity level versus number of trucks used is shown in Figure 4.9.

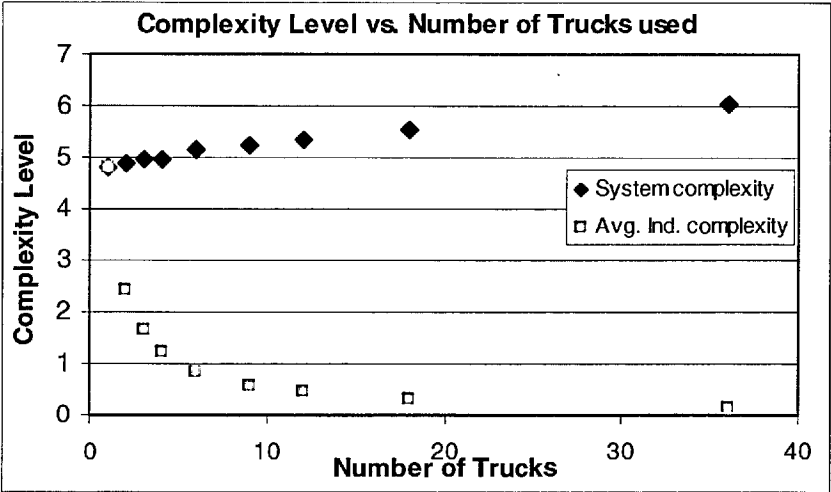


Figure 4.9: Complexity Level vs. Number of Trucks in fixed route vehicle routing

In general, given the same number of trucks used, fixed routes will be less complex than flexible routes. However, from Table 4.6 and Figure 4.9, when the number of trucks increases, the complexity level of the system increases as well. It is also worth noting that although the system complexity increases as the number of trucks increases, the average complexity level of individual trucks or routes actually decreases because each truck will visit fewer customers.

4.8 Degree Distribution in Vehicle Routing Network

The paths or sequences which vehicles take within the logistics configuration create a vehicle routing network. Figure 4.3 illustrates an example of a centralised vehicle routing network. In this section, we analyse vehicle routing networks under a framework of graph theory. There are two types of path between nodes in a vehicle routing network; the one that directs toward a node and the one that directs away from a node. When, for example, there is a path direct from node A toward node B (a vehicle travels from A to B) and a path direct from node B toward node A (a vehicle travels from B to A) as shown in Figure 4.10, we will count the number of edges for each A and B node as two.

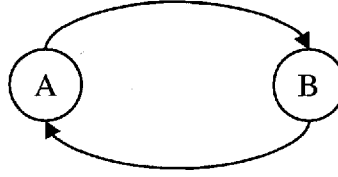


Figure 4.10: Illustration of edges direct toward and direct away from nodes

With the convention described above, we count number of edges attached to each node. From the entropic complexity formula, we know that there will be a journey from i to j if and only if $p_{ij} > 0$. Let S_{ij} be a binary variable represents whether a path from node i to node j exists (0 if there is no path, 1, if a path exists). Therefore, we can express S_{ij} as

$$S_{ij} = \begin{cases} 1 & \text{if } p_{ij} > 0 \\ 0 & \text{if } p_{ij} = 0 \end{cases}$$

The total number of edges attached to node i (S_i) is number of edges from i toward other nodes plus number of edges from other nodes toward i , which can then be expressed as

$$S_i = \sum_{j=1}^n S_{ij} + \sum_{j=1}^n S_{ji}, \quad \forall i$$

There is no double counting here because S_{ii} is always zero. Using the above formula to obtain all S_i , we can calculate $p(k)$ by dividing the number of nodes with $S_i = k$ by the total number of nodes. Therefore, we now can plot k against $p(k)$ for the vehicle routing networks to see what

distribution they follow (e.g. the power-law distribution discussed in the previous section). We present results from two different planning horizons: 5 days (typical planning horizon in practice) and 6000 days (steady state) of each logistics configuration and each vehicle size.

4.8.1 Probability Distribution Plots

The probability distributions of degree of vertex are plotted in Figures 4.11 to 4.13 for the first demand profile and in Figures 4.14 to 4.16 for the second demand profile. A regression line is shown in each plot as the best fit line; the equation of regression line with the exponent term and R-squared are also displayed. The R-squared is a coefficient of determinant indicating goodness of fit between the regression line and the actual plot. The higher the R-squared is, the better the fit. In general, an R-squared value of 0.4 or above is considered to show some consistency between the regression line and the data.

Probability distribution plots are shown for 4 logistics configurations which are configuration A (1 hub), C (2 hubs), E (3 hubs), and G (3 hubs), respectively. Each logistics configuration is selected to represent a different number of warehouses in the logistical network. However, we have two configurations with 3 warehouses (E and G). Configuration G is added because it is the least centralised configuration. Each figure below shows distribution plots for large trucks, medium trucks, and small trucks respectively. In each figure, each plot labelled by (a), (b), (c), (d), (e), (f), (g), and (h) displays degree distribution of each condition (i.e. logistics configuration, planning horizon).

4.8.1.1 Degree Distribution Plots for First Demand Profile

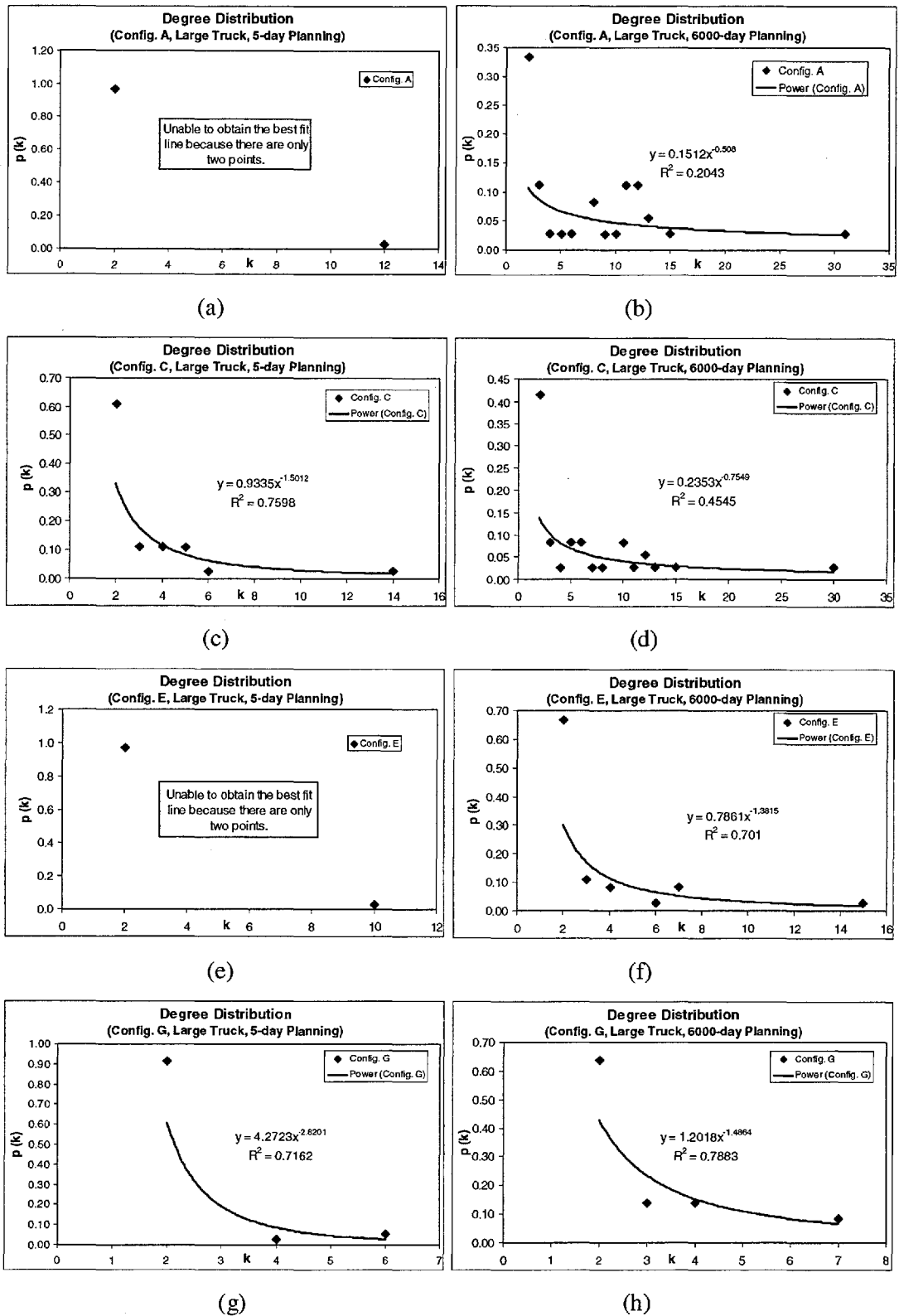
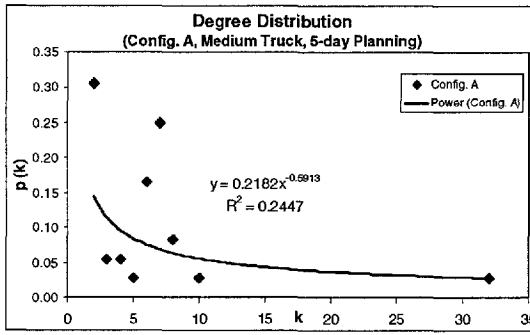
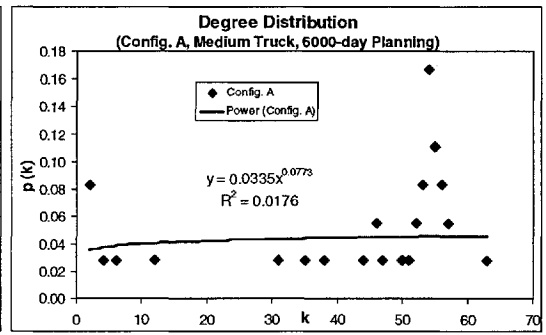


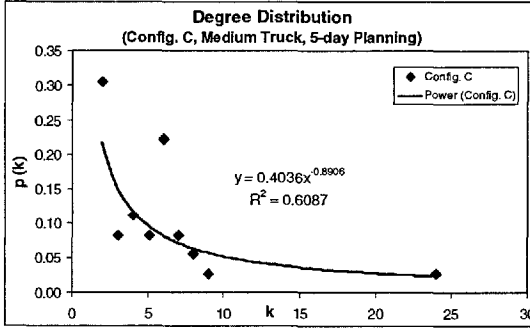
Figure 4.11: Degree distribution plots of degree of vertex for large truck (1st demand profile)



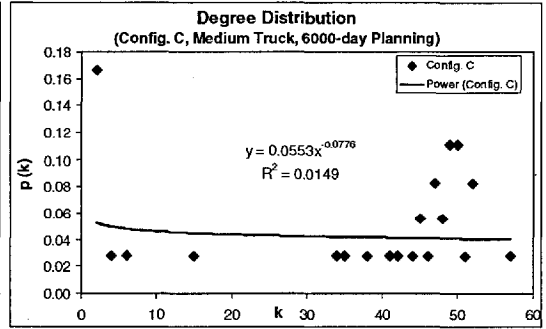
(a)



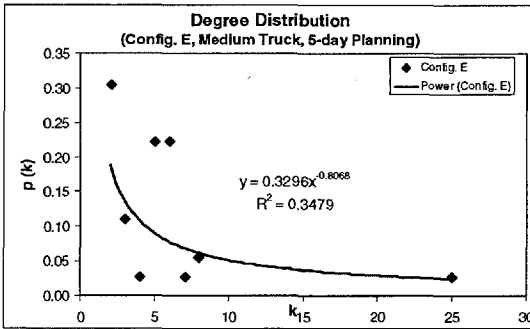
(b)



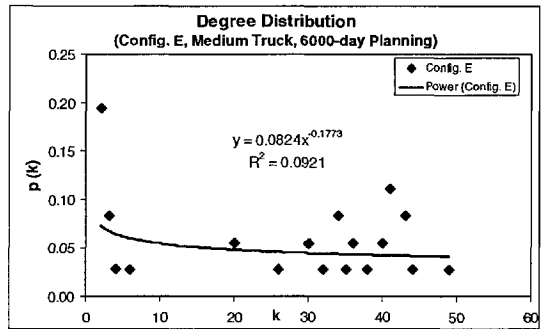
(c)



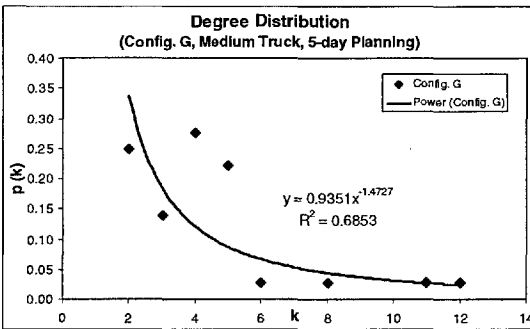
(d)



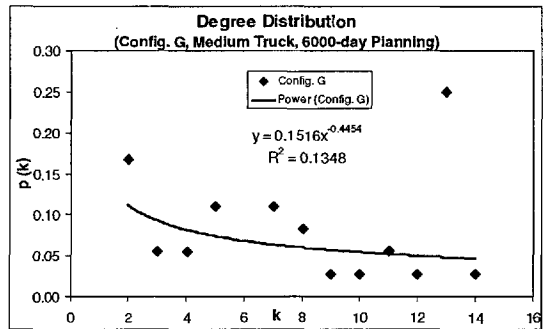
(e)



(f)

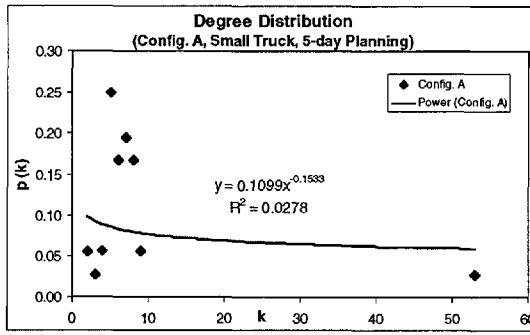


(g)

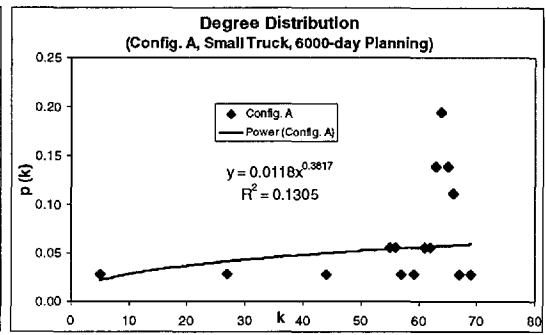


(h)

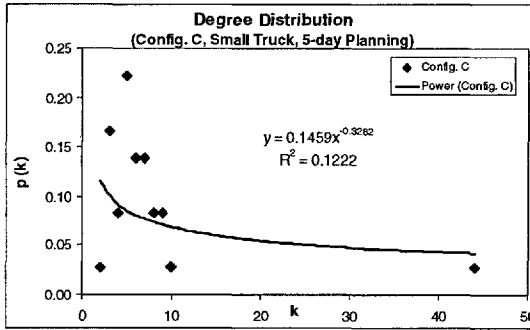
Figure 4.12: Degree distribution plots of degree of vertex for medium truck (1st demand profile)



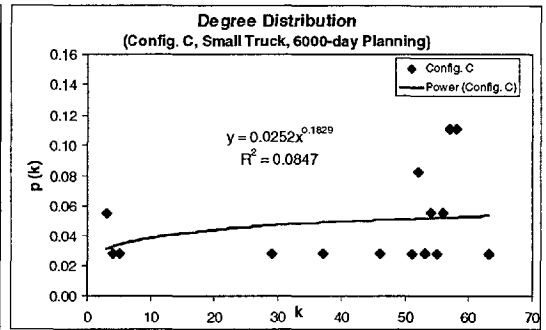
(a)



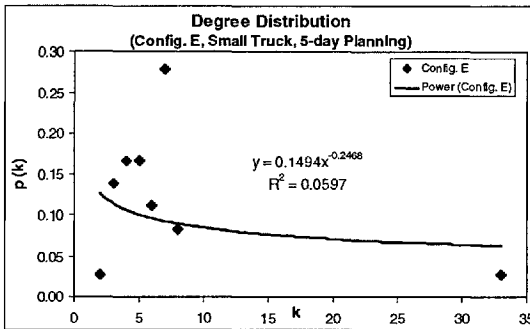
(b)



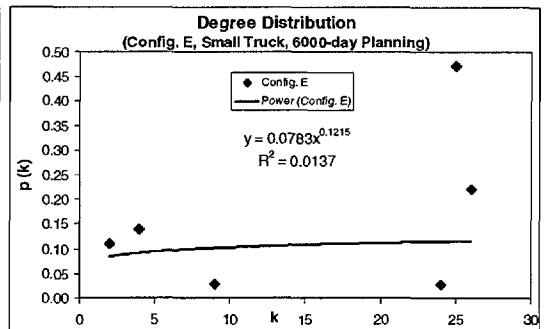
(c)



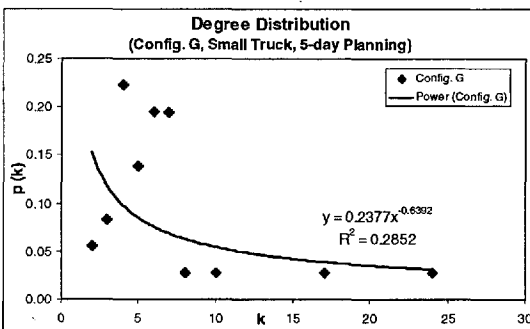
(d)



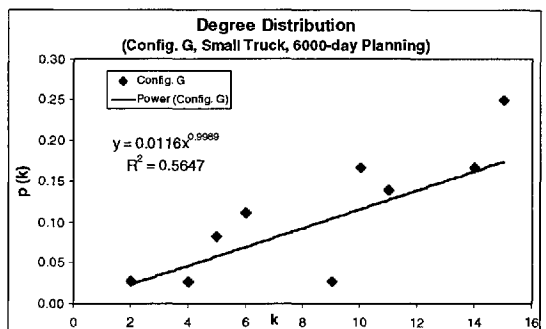
(e)



(f)



(g)



(h)

Figure 4.13: Degree distribution plots of degree of vertex for small truck (1st demand profile)

4.8.1.2 Degree Distribution Plots for Second Demand Profile

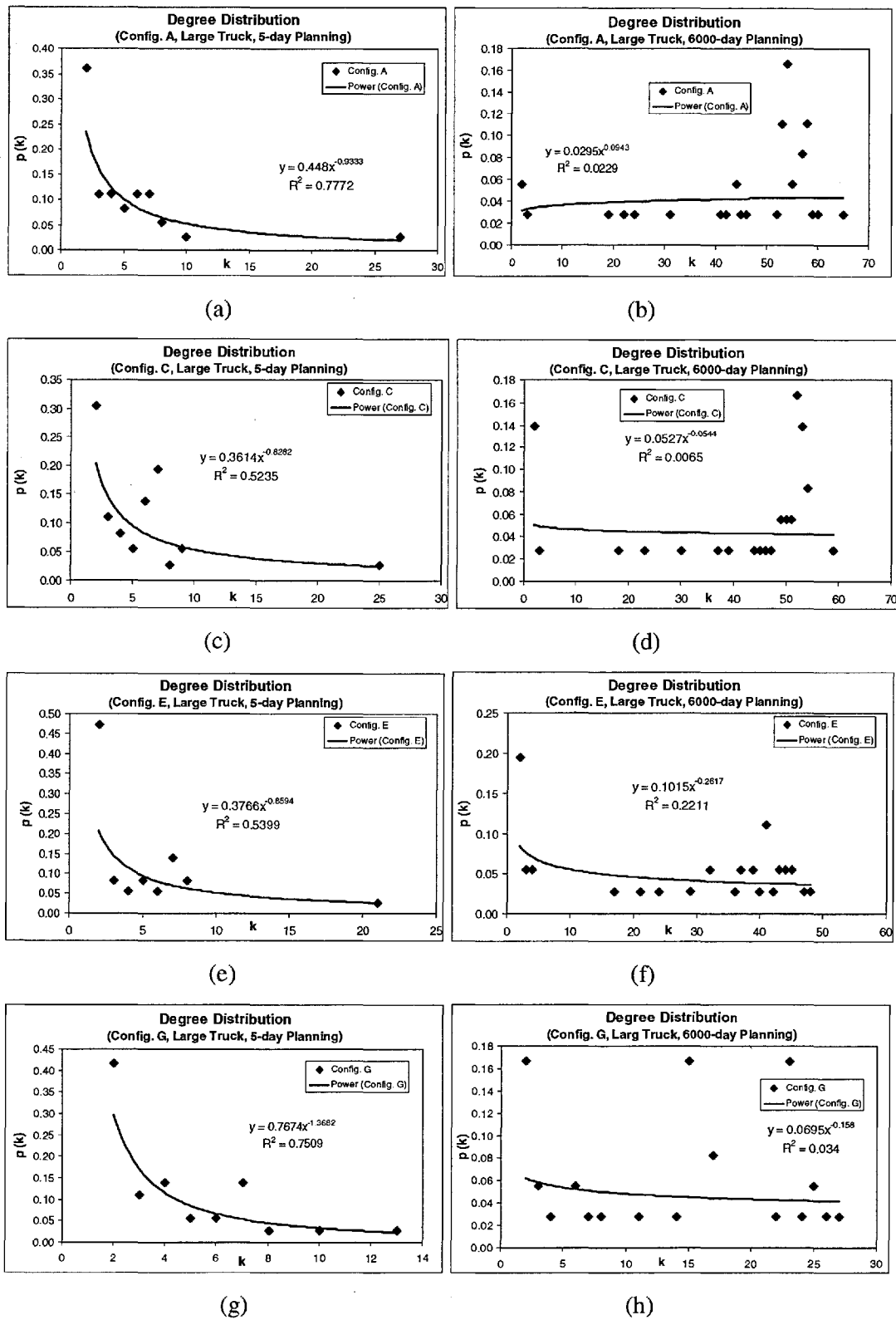
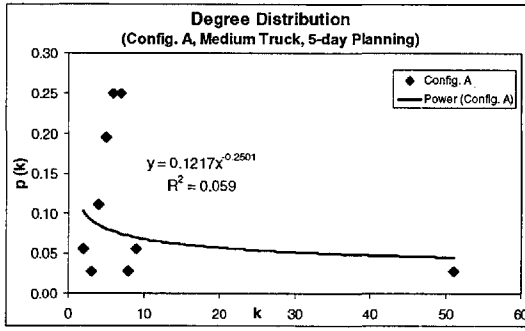
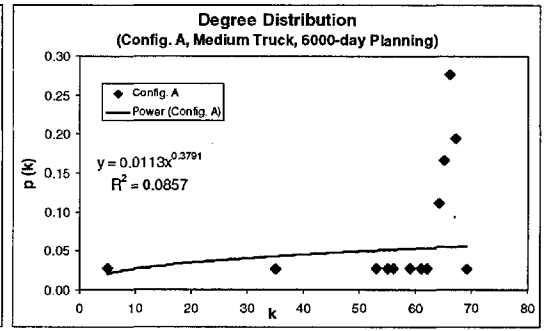


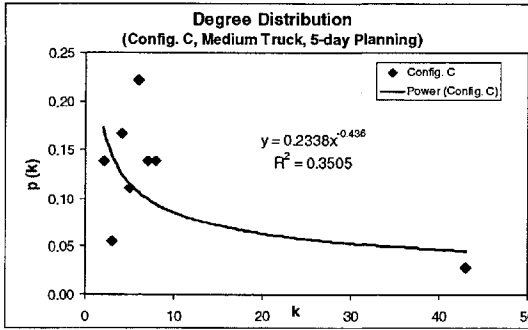
Figure 4.14: Degree distribution plots of degree of vertex for large truck (2nd demand profile)



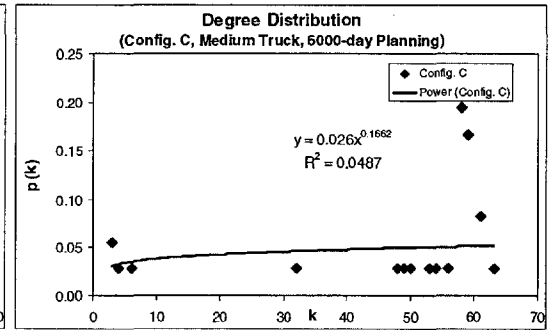
(a)



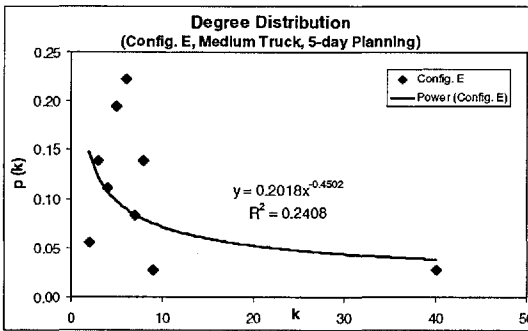
(b)



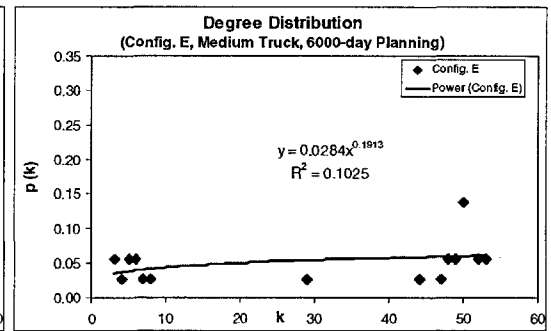
(c)



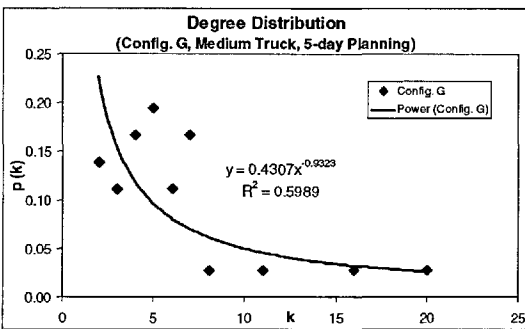
(d)



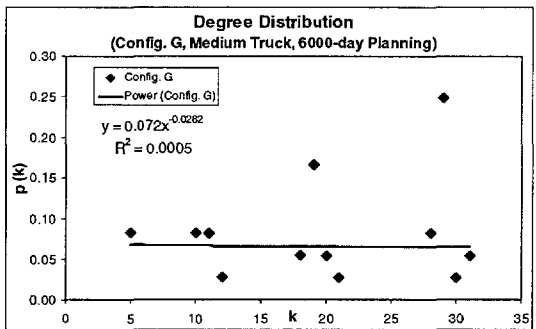
(e)



(f)



(g)



(h)

Figure 4.15: Degree distribution plots of degree of vertex for medium truck (2nd demand profile)

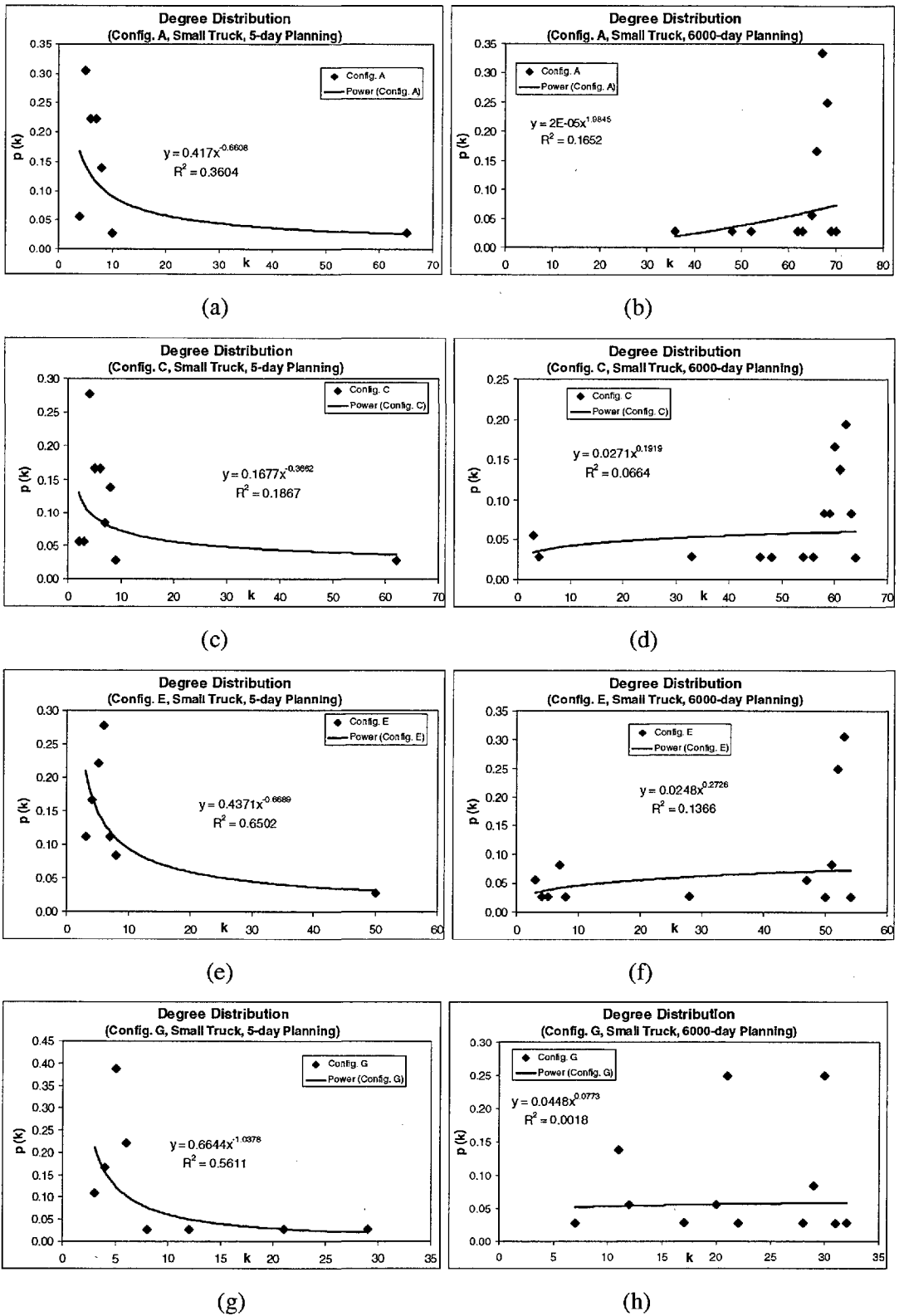


Figure 4.16: Degree distribution plots of degree of vertex for small truck (2nd demand profile)

From all the plots above, for both demand profiles, there are a few things we can observe. Though not all the plots follow the power-law closely, there are some vehicle routing networks that exhibit a high consistency with the power law (i.e. negative exponent power distribution

function with good R-squared). These networks tend to come from the configuration where the planning horizon is short and a large vehicle size is used. We summarise this finding in a matrix classifying level of consistency with the power-law degree distribution of vehicle routing networks as shown in Figure 4.17.

		Vehicle Size	
		Small	Large
Planning Horizon	Short	Moderate Consistency	High Consistency
	Long	Low Consistency	Moderate Consistency

Figure 4.17: Level of consistency with the power-law degree distribution

It also seems that the logistics configuration has some influence on network structure, in particular the consistency with the power law. This is probably because a network topology with multiple hubs is more consistent with the power law; resulting in a relatively scale-free network. Such networks are likely to be more robust (i.e. loss of one hub may not disable the network).

4.9 Statistical measures of degree of vertex

From various probability distribution plots in Figures 4.11 to 4.16, we can obtain statistical measures of degree of vertex in the vehicle routing network such as mean and variance. Mean and variance of degree of vertex represent estimation of expectation and variability of degree of vertex. They can hold some economic implications for the vehicle routing network. A greater a mean and variance can translate into a greater effort for planning and coordination. The mean and variance can be calculated as follows;

Mean or Average of degree of vertex:

$$\bar{k} = \sum_{k=1}^{\infty} p(k)k \quad (4.2)$$

Variance of degree of vertex:

$$\sigma_k^2 = \sum_{k=1}^{\infty} p(k)(k - \bar{k})^2 \quad (4.3)$$

For the first demand profile, the results are shown in Table 4.7. The mean and variance in Table 4.7 are then categorised and plotted in Figure 4.18. Table 4.8 shows results from the second demand profile. Column plots of mean and variance from the second profile are displayed in Figure 4.19.

First Demand Profile

Vehicle Type	Configuration	5-Day Planning		6000-Day Planning	
		Mean (k_{bar})	Var (σ_k^2)	Mean (k_{bar})	Var (σ_k^2)
Large Truck	A	2.28	2.70	7.17	35.25
	C	3.11	4.77	5.94	31.39
	E	2.22	1.73	3.17	6.31
	G	2.28	0.92	2.83	2.08
Medium Truck	A	5.72	25.37	43.56	342.41
	C	5.00	14.72	36.22	367.78
	E	4.83	15.14	25.11	284.54
	G	4.17	5.03	12.94	62.55
Small Truck	A	7.33	62.56	59.67	139.56
	C	6.72	43.76	49.44	301.41
	E	6.17	23.36	35.56	367.36
	G	6.11	15.77	19.56	67.58

Table 4.7: Mean and Variance of degree of vertex (1st demand profile)

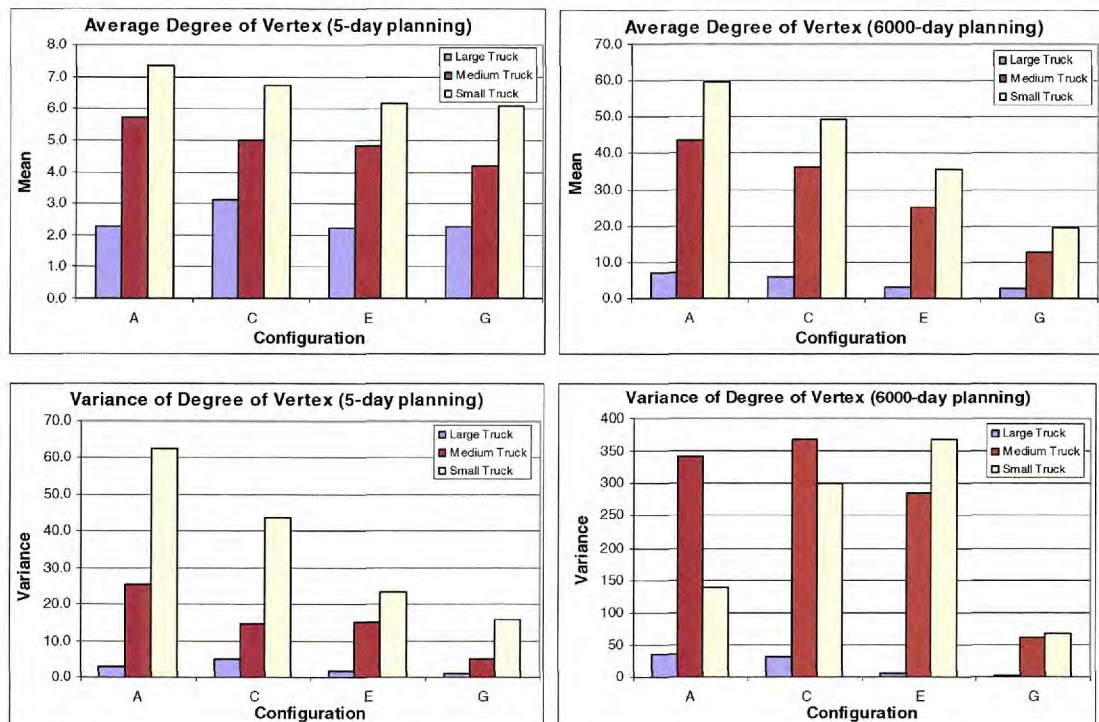


Figure 4.18: Mean and Variance of degree of vertex (1st demand profile)

Second Demand Profile

Vehicle Type	Configuration	5-Day Planning		6000-Day Planning	
		Mean (k_{bar})	Var (σ_k^2)	Mean (k_{bar})	Var (σ_k^2)
Large Truck	A	4.83	18.97	45.83	284.08
	C	5.17	16.42	40.00	357.22
	E	4.39	12.74	27.11	308.60
	G	4.17	6.97	13.94	74.05
Medium Truck	A	7.00	57.89	61.94	128.00
	C	6.28	42.15	51.22	307.56
	E	6.33	35.67	38.39	381.90
	G	5.61	13.13	20.50	70.86
Small Truck	A	7.83	95.08	65.06	41.33
	C	6.89	89.88	54.17	267.25
	E	6.61	55.74	40.94	377.05
	G	6.22	24.06	22.28	55.81

Table 4.8: Mean and Variance of degree of vertex (2nd demand profile)

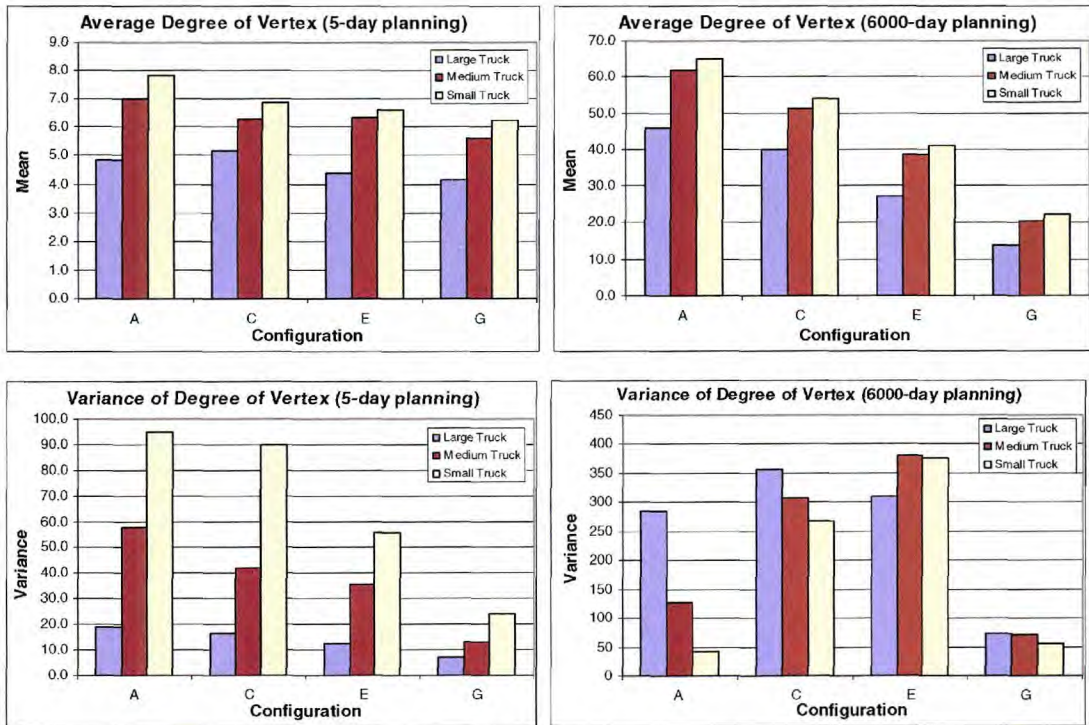


Figure 4.19: Mean and Variance of degree of vertex (2nd demand profile)

Some observations can be made here. In both profiles, the average degree of vertex tends to be the highest in configuration A (centralised distribution). Only when large vehicles are used and the planning horizon is short, does configuration C have the highest average degree of vertex.

In all other conditions the average degree of vertex in each network has the following trend:

$$\text{Average degree of vertex of configuration } A > C > E > G$$

The trend is even more obvious when the planning horizon is long (i.e. in steady state). That configuration G gives the lowest average degree of vertex even though it has the same number of hubs (3) as configuration E is probably because customer zones are assigned to each warehouse more evenly (see Table 4.1). When the planning horizon is short, the variance of degree of vertex also exhibits the same trend as above. However, in long planning horizons this trend does not hold true. It is also worth noting that vehicle size (load capacity) greatly impacts the average degree of vertex all cases. The larger the vehicle size, the lower the average degree of vertex.

4.10 Managerial Implications

The results from computational experiments based on our case study show that the choice of logistics configuration, vehicle capacity, and demand profile affects the level of complexity in vehicle routing. However, if supply chain managers are concerned about long-term complexity rather than short-term complexity, they should look at steady state complexity. Although it is often believed that centralised distribution is simpler and less complex, based on our results it is revealing that centralised distribution is actually potentially more complex in terms of vehicle routing. As traffic conditions are often difficult to predict, having a low level of vehicle routing complexity (i.e. repeating a small number of routes) should be more desirable to maintain reliable and on time delivery. Therefore, if vehicle routing is part of the company's strategic advantage, the supply chain manager may consider decentralised distribution (e.g. having more warehouses) to reduce vehicle routing complexity. However, it may be difficult and not cost effective to change logistics configurations. Our results also show that increasing vehicle capacity to certain levels can also provide a dramatic effect in reducing complexity levels, especially over long horizons. Although degree distribution of networks such as those of vehicle routing is of interest among network researchers, average degree of vertex may be of interest from a managerial point of view as its pattern versus logistics configuration is similar to that of complexity level.

4.11 Summary and Discussion

Supply chain management has many objectives which often conflict with each other. Nevertheless, it is imperative that supply chain managers understand all aspects of their supply chain performance. Vehicle routing can offer a strategic advantage to a company through reliable and on time delivery at minimum cost. As complexity in vehicle routing can be regarded as a potential obstacle to achieving such performance, a supply chain manager should also take vehicle routing complexity into account when making supply chain decisions. In this chapter we propose an approach to evaluate vehicle routing complexity. Our approach can reveal insights regarding complexity levels over different choices of logistics arrangements (e.g. logistics configuration, vehicle capacity, etc.) as we demonstrate in our case study. Our results also suggest that decentralised distribution and large vehicle capacities can both reduce vehicle routing complexity. However, a simulation should be performed to confirm an appropriate level of capacity increase, as we see in the second demand profile case that some level of capacity increase may result in a higher complexity level (e.g. complexity level of medium trucks is higher than that of small trucks when planning horizons are long). The results also show that although in the short term complexity level may be low, it can be much higher when we consider the long term. However, there exists a steady state complexity which is not exceeded even for long planning horizons. In fixed route vehicle routing, the system complexity increases as the number of trucks used increases. However, the average complexity contributed by each route decreases when more trucks are used.

In this chapter we also investigated topological properties of vehicle routing networks. The vehicle routing networks in our case study exhibit both power-law and non power-law degree distributions. We also discover that the pattern of *average degree of vertex vs. logistics configuration* tracks that of complexity level.

Chapter 5

Multi-Agent Based Supply Chain Modelling and Operational Complexity

Supply chain systems are often characterised as non-linear dynamic systems where certain information is global or known to all members in the system, while some other information is local and known only to individual or neighbouring members. Therefore, analytical methods or equation based methods may struggle to capture the exact characteristics of the system. Furthermore, the discrete nature of the supply chain makes it even more difficult to evaluate supply chain system analytically. Nonetheless, it is important that a supply chain manager understands well and is able to predict his or her supply chain behaviours, in particular supply chain efficiency, complexity, and service levels, as they are important performance measures in supply chain operations. It is widely documented that an IT revolution has taken place in this sector. We are interested in whether simulation can be used to assess the benefits of potentially more complex transactional systems. In this chapter, we propose assessment of supply chain performance through a simulation method based on multi-agent based modelling concept. We develop a modelling framework, define supply chain policies, and create a simulation model on the NetLogo modelling platform. We then perform simulation on our hypothetical case study to evaluate supply chain performance in a supply chain environment under different conditions of operational complexity.

5.1. Introduction

Supply chain management has frequently been cited as having a central role for success in many production and service organisations in the past decade. Supply chain management strategies,

however, have continually evolved as globalisation and competition become more intense, and as the supply chain environment changes with developments such as deregulation, emerging opportunities in new markets, etc. Business entities within the supply chain nowadays are becoming more and more distributed, not only in terms of the distance between them but also in terms of their decision making processes. Nonetheless, every member in the supply chain still needs to interact and is interdependent with other members in the supply chain. To gain valid insights from supply chain systems analysis, it is essential that the tools used are able to capture and characterise the discrete, dynamic, and distributed nature of such supply chain systems.

Simulation is often used as a supply chain management decision support tool when there is a need to incorporate dynamic behaviour and uncertainty into the analysis (Shapiro, 2001; van der Zee and van der Vorst, 2005). Application of multi-agent based modelling and simulation in supply chain management has been a natural approach since it can characterise well the nature of supply chain dynamics (Swaminathan *et al.*, 1998). There are a number of research outputs from multi-agent based supply chain simulation for various objectives and various aspects of supply chain management. Caridi and Cavalieri, (2004) provided an extensive review of multi-agent systems applied to production planning and control. However, the main focus is to improve supply chain performance measures such as inventory (Ahn and Park 2003), profitability (Allwood and Lee, 2005), quality of service (Davidsson and Wernstedt, 2004), etc.

Multi-agent based modelling techniques can allow agents to explicitly communicate with each other by facilitating communication protocols such as knowledge query message language (KQML) (García-Flores and Wang, 2002; Gjerdrum *et al.*, 2001b). Multi-agent based modelling techniques are applicable to virtually all aspects of supply chain decision-making process such as third party logistics (Ying and Dayong, 2005), vehicle planning (Fischer and Gehring, 2005), resource allocation (Kaihara, 2003), production planning and scheduling (Choi *et al.*, 2004; Pěchouček *et al.*, 2004), supplier-buyer selection (Lee *et al.*, 2004), procurement (Julka *et al.*, 2002a, 2002b), order fulfilment (Lin *et al.*, 1998; Lin and Shaw, 1998; Lin and Pai, 2000), etc. Multi-agent based simulation techniques can also let us observe *emergent behaviour* of the systems. An emergent behaviour is a complex system behaviour as a result of local interactions among simple agents in the system. One of the most well-known emergent behaviours in the supply chain system is the 'Bullwhip Effect' (Lee *et al.*, 1997).

Supply chain systems nowadays are increasingly complex, and there are many proposals to characterise and quantify complexity of supply chain systems in various aspects, such as supply chain network complexity (Meepetchdee and Shah, 2007a), product flow and stock complexity (Wu *et al.*, 2002), manufacturing system complexity (Funk, 1995; Calinescu *et al.*, 1998; Yu

and Efstathiou, 2002), manufacturing software systems complexity (Phukan *et al.*, 2005), etc. Recently, investment in supply chain IT systems has increased dramatically in all industrial sectors. This provides more options for supply chain operations but will increase transactional complexity as well. Though complexity is usually regarded as undesirable, we shall explore the potential benefit of such complexity. Along with efficiency and complexity, fill rate is another very important supply chain performance measure because the fill rate indicates how well the demand is being served. In this chapter, we characterise and model a supply chain system based on a multi-agent based modelling framework and perform simulation on a hypothetical case. We also examine the effect of transactional complexity (via different ordering options) on supply chain performance and explore relationships between efficiency, complexity, and fill rate in supply chain operations.

5.2. Modelling framework and agents for supply chain system

A supply chain is a customer-demand driven system made up of several entities interacting with each other. Since we will adopt the multi-agent paradigm for modelling and analysis of the supply chain, from now on we will refer to each entity in the supply chain as an agent. These agents interact with other agents based on their internal pre-specified policies. The supply chain network structure dictates which agents will have direct interaction with which other agents. Each agent has its own internal procedure to arrive at its action or decision. The action or decision of each agent depends on its pre-defined procedure, its internal variables and resources, and information from its neighbouring agents. Furthermore, the action of each agent is either triggered by its internal schedule or information from its neighbouring agents. There are two categories of agents in this framework: *structural agent* and *logical agent*. The structural agent is the agent that structurally holds supply chain management policies (e.g. inventory control, production planning, etc.), while the logical agent is the agent that does not hold any policy but is logically necessary in the system for the completeness of the model. The structural agents in our framework are Retailer agent, Warehouse agent, and Manufacturer agent. The market agent and Raw Material agent, in this framework, are logical agents. We define general characteristics of agents in our supply chain system as follows:

Market Agent: A market agent is a logical agent which generates demand of product for each retailer in each period. It is considered as a collection of customers who demand the product. The demand can be constant, fixed (e.g. from empirical data, etc.), or stochastically generated according to pre-specified distributions. The demand from the market will initiate activities within the supply chain system (customer-demand driven system).

Retailer Agent: A retailer agent serves customer demand (from the market agent) by holding some product inventory based on its inventory management policy. If the product is available in the inventory at the time the customer demand arrives, the demand will be fulfilled immediately and the retailer's inventory will be updated. However, if the product is not available or if the product inventory is insufficient to fulfill customer demand in full, the retailer will either record a backorder so that unfilled demand will be served as soon as the product becomes available in future periods, or the sales are lost. Whenever the retailer's inventory level reaches or falls below its re-order level, it will request product delivery from a warehouse. The quantity of product to be ordered depends on the retailer's order policy, which in turn depends on the complexity of its transactional system.

Warehouse Agent: A warehouse agent holds product inventory to be distributed to retailers upon request, but there may be some processing and delivery lead time. After logging delivery requests from the retailers, the warehouse records the orders from all retailers and plans deliveries according to the lead time. If the inventory in the warehouse is not sufficient to serve replenishment orders from all retailers in full, the warehouse will deliver as much product as possible subject to delivery constraints (e.g. minimum order quantity, fixed order increment, etc.). When the warehouse inventory is not enough to serve demand from all retailers for any particular delivery period, it will determine the quantity of the product to be distributed partially to each retailer according to its pre-defined rules. The warehouse manages its inventory according to its policy and requests replenishment from a manufacturer when needed; again the scope for these transactions can be more or less complex.

Manufacturer Agent: A manufacturer receives replenishment requests from the warehouse. The manufacturer may or may not hold finished goods inventory, depending on its policy (e.g. make to stock, make to order, etc.). The manufacturer will schedule its production plan according to replenishment requirements from the warehouse, production capacity, flexibility in production schedule change, finished goods inventory level, and raw material availability, etc. Once the manufacturer updates its production schedule and pulls raw material into the production process, the raw material inventory level will be updated. If outputs from the production schedule and finished goods inventory cannot serve replenishment requests in full in any particular period, the manufacturer will deliver as much product as possible subject to delivery constraints (e.g. minimum order quantity, fixed order increment, etc.). When necessary the manufacturer will request raw material from the raw material agent according to its raw material inventory management policy.

Raw Material Agent: A raw material agent responds to material requests from the manufacturer. In this framework, the raw material agent is a logical agent that merely supplies material to the manufacturer upon request. There may be some lead time in raw material delivery, however, there are neither any elaborated supply chain management activities within the raw material agent nor any delivery constraints between the manufacturer and the raw material provider.

5.3. Supply chain system performance measures

5.3.1 Supply Chain Efficiency

At the operational level of supply chain management, inventory is a critical source of supply chain strategic advantage, especially for perishable products or products that can become obsolete. Major costs incurred by inventory are storage cost, perishing or obsolescence cost, handling cost, etc. Inventory turnover ratio is a measure of the efficiency with which inventory is managed; computed by dividing the amount of goods sold by average inventory for a period. Thus, having a high inventory turnover ratio means that the cost of the inventory can be recovered more quickly, meaning more efficient inventory management. In this chapter, we use inventory turnover ratio as our measure of supply chain inventory efficiency. The inventory turnover ratio can be calculated as follows:

$$\text{Inventory Turnover Ratio} = \frac{\text{Total (Annual, Monthly, etc.) Sales}}{\text{Average Inventory Level}} \quad (5.1)$$

Another measure of efficiency is the average number of orders in each evaluation interval. The average number of orders is a measure of efficiency in ordering transaction and delivery. The lower the average number of orders, the more efficient the transaction and delivery since each order will incur a transaction cost, order handling cost, logistics cost, etc.

$$\text{Average number of orders} = \frac{1}{N} \sum_{i=1}^N R_i \quad (5.2)$$

where N is the total number of simulation runs and R_i is the number of orders in the i^{th} run.

5.3.2 Supply Chain Complexity

There are various types of complexity in any supply chain system. In this chapter, we will examine complexity in supply chain operations because a high complexity in supply chain

operations may imply that more resources and effort may be required to synchronise and coordinate activities within the system and a more complex IT system must be employed. Fluctuation in inventory levels frequently makes it more burdensome to manage inventory, since it will be more difficult to anticipate resources required to manage, store, and handle inventory. Mean Absolute Percentage Error (MAPE) is usually used to measure forecasting accuracy (Nahmias, 2005). However, when applied to inventory level, it can also reflect the magnitude of fluctuation in the inventory. To represent operational complexity in inventory management, we propose a modified Mean Absolute Percentage Error which uses average sales as a denominator as a measure of inventory complexity. We use *average sales* instead of average inventory because otherwise the MAPE can be distorted by the inventory management policy (e.g. holding very high or very low inventory level). The modified Mean Absolute Percentage Error can be calculated as follows;

$$MAPE = \left[\frac{1}{n} \sum_{j=1}^n \frac{|\text{Inventory Level}_{\text{period } j} - \text{Average Inventory Level}|}{\text{Average Sales}} \right] \times 100 \quad (5.3)$$

where n is the total number of period in consideration in each evaluation interval.

To characterise the complexity of ordering and delivery transactions, we propose a complexity measurement based on entropy concepts. This method defines complexity in terms of the expected amount of information required to describe state of the system. We will apply this concept to different delivery sizes in the supply chain as follows:

$$\text{Entropic Complexity} = - \sum_{k \in D} p_k \log_2 p_k \quad (5.4)$$

where D is a set of all possible delivery sizes and p_k is the probability of having delivery size k .

From the formula given above, the system will become more complex when there are many different delivery sizes and when the chance (probability) of having any delivery size is disperse (wide probability distribution), which in turn will be reflected in a more complex IT system and order handling process. The complexity will be reduced if there are only few delivery sizes or if the probabilities of having a few particular delivery sizes are much greater than others (i.e. a tight probability distribution).

We can consider this complexity as a complexity in transactional operations such as IT systems of purchasing, ordering, or product receiving functions. When there are many different order sizes or delivery sizes, the IT systems that support such operations will have to be more

complex since there are many possible scenarios that can occur in the system. On the other hand, the increased degrees of freedom may result in improved efficiency. The complexity is accentuated by scattering probability distribution of the possible scenarios. The maximum complexity for any set of possible scenarios occurs when all scenarios are equiprobable (Frizelle and Woodcock, 1995); while minimum complexity occurs when there is only one possible scenario, where the complexity becomes zero.

5.3.3 Fill Rate

The primary function of a supply chain system is to fulfil customer demand. Ideally this demand should be fulfilled 100% at all time. However, this rarely happens in reality because such supply chain systems would have to hold a huge amount of inventory or have to quote a very long lead time. In either case, the inventory cost may not justify 100% fulfilment or the long lead time may turn the customers away to other competitors. Nevertheless, maintaining a high fill rate is important in supply chain management because it is directly linked to customer satisfaction. Here, we define the fill rate as follows:

$$\text{Fill Rate} = \frac{\text{Demand that is satisfied from inventory or within promised lead time}}{\text{Total Demand}} \quad (5.5)$$

Given these characteristics, we develop a simulation model, in particular to identify the benefit of complexity in terms of the performance measures of inventory and service level.

5.4. Implementation

5.4.1 Simulation Platform

Based on our framework, we implement the multi-agent supply chain model on the NetLogo modelling platform (Wilensky, 1999). NetLogo is a programmable multi-agent modelling environment for simulating complex systems in which modellers can give instructions to a large number of independent “agents” to operate concurrently. NetLogo also provides a user-friendly programming and graphical interface which allows modellers to observe the performance of individual agents or macro-level patterns that emerge from the interaction of many agents in the system. NetLogo also allows modellers to extend its capability by writing new commands and reporters in *Java* and use them in their models. Therefore, NetLogo is both a capable and flexible platform that we can use to develop our supply chain model. In our model, we do not

need a formal communication protocol between agents. The coordination in the supply chain system takes place by an implicit mechanism where the behaviour logic of the single agents are axante defined and known (Keilmann, 1995).

5.4.2 Supply Chain Model

As discuss in the modelling framework, our supply chain system consists of a market, retailers, warehouse, manufacturer, and raw material provider. In our model, we consider a supply chain system where there are 4 retailers, 1 warehouse, and 1 manufacturer. The market can be considered as a collection of customers, while the raw material provider can be considered as a collection of material suppliers as displayed in Figure 5.1.

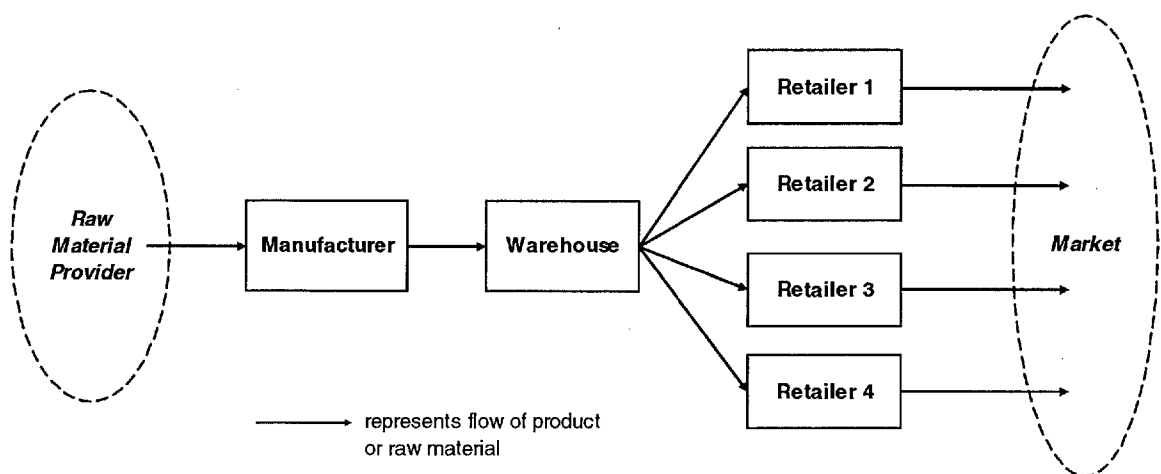


Figure 5.1: Supply Chain Network Structure

5.4.3 Supply Chain Management Policy

In this model, the market is a logical agent and does not implement any supply chain management policy. It merely generates demand according to pre-specified formula or distribution (constant, fixed, or stochastic). The demand generated by the market activates the retailers' activities to fulfil customer demand and, when applicable, to re-order product from warehouse. The orders from retailers to the warehouse must be at least equal to minimum order quantity (RMOQ) and must increase in steps of fixed order increment (RFOI). In our model the retailers adopt an (s, S) policy for their inventory management (Nahmias, 2005).

If u is the inventory level in any period, then

Order $\lceil S - u \rceil$, if $u \leq s$

Do not order, if $u > s$

$\lceil S - u \rceil$ denotes round up to the closest $RMQ + Int \times RFOI$ that is equal or greater than $S - u$ where Int is a non-negative integer.

The activities of a market and a retailer and their interaction are shown in Figure 5.2.

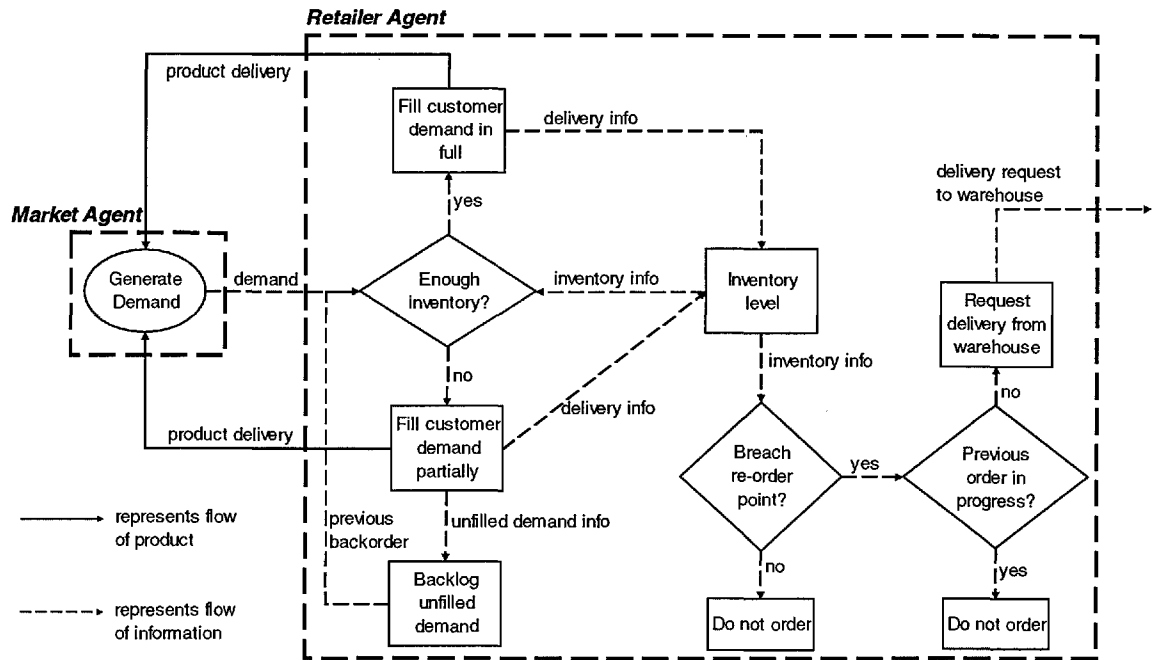


Figure 5.2: Market agent and Retailer agent's supply chain policy

A warehouse receives product orders from retailers in the downstream supply chain. The warehouse may require some time to process the order, to perform packaging, and to deliver the product to retailers. This lead time is supposed to be agreed by the warehouse and the retailers. Hence, the warehouse will schedule delivery according to this lead time. The warehouse will hold product inventory and will schedule replenishment requests from the manufacturer as per its inventory management policy. In our model, the warehouse adopts the Distribution Requirement Planning (DRP) scheme for its inventory management by maintaining a certain level of inventory as a based stock. The warehouse can project its future inventory level since it can deduce future inventory levels from its current inventory level and scheduled deliveries to retailers (according to lead time). The warehouse then creates a replenishment plan for the

manufacturer according to the expected future inventory level and its base stock target. The replenishment request from warehouse to manufacturer must also exceed a minimum order quantity (WMOQ) and increase in steps of a fixed order increment (WFOI). However, WMOQ and WFOI can be different from RMOQ and RFOI. Figure 5.3 shows an example of warehouse inventory profile when starting inventory (I_0) equals based stock target (B_T) and how it can create a replenishment plan for the manufacturer. The table on the figure summarises replenishment requests for each period. The symbol $\lceil \cdot \rceil$ denotes round up to the closest $WMOQ + Int \times WFOI$ that is equal or greater than an expression inside $\lceil \cdot \rceil$ where Int is a non-negative integer.

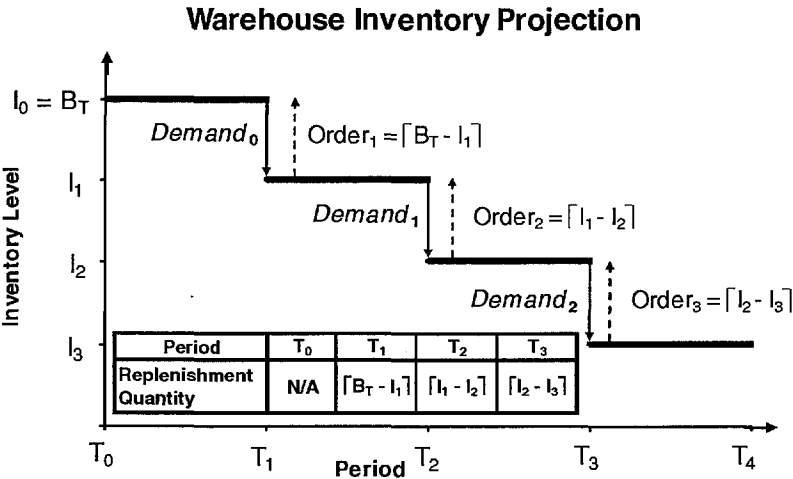


Figure 5.3: Example of Warehouse Inventory Projection when $I_0 = B_T$

We can generalise the replenishment plan where starting inventory (I_0) may or may not equal B_T as shown in the Table 5.1.

Period	T ₀	T ₁	T ₂	T ₃
Inventory Level	I_0	I_1	I_2	I_3
Replenishment Quantity	N/A	0, if $I_1 \geq B_T$ $\lceil B_T - I_1 \rceil$, if $I_1 < B_T$	0, if $I_2 \geq B_T$ $\lceil I_1 - I_2 \rceil$, if $I_2 < B_T$ and $I_1 < B_T$ $\lceil B_T - I_2 \rceil$, if $I_2 < B_T$ and $I_1 \geq B_T$	0, if $I_3 \geq B_T$ $\lceil I_2 - I_3 \rceil$, if $I_3 < B_T$ and $I_2 < B_T$ $\lceil B_T - I_3 \rceil$, if $I_3 < B_T$ and $I_2 \geq B_T$

Table 5.1: Generalised replenishment quantity for warehouse’s replenishment plan

The procedure within the warehouse and its interaction with other agents is shown in Figure 5.4.

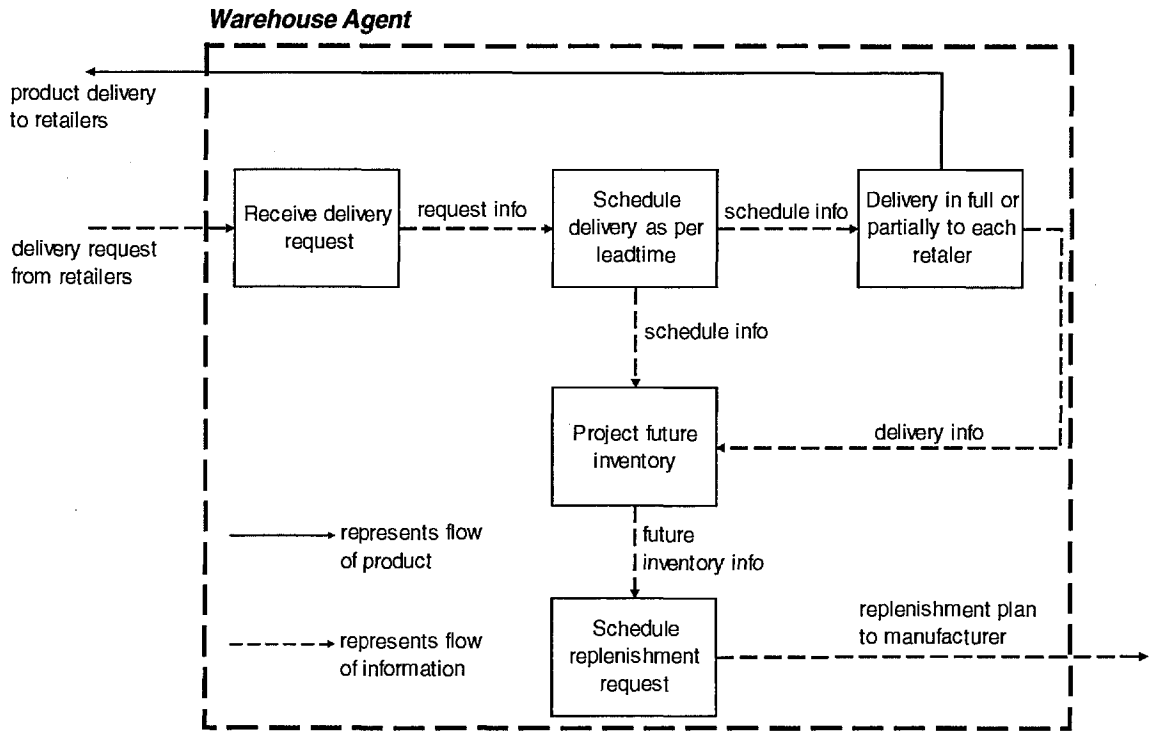


Figure 5.4: Warehouse agent's supply chain policy

A manufacturer receives a replenishment plan from the warehouse and schedules its production to fulfil the warehouse's replenishment plan. In our model, the manufacturer adopts a make-to-order policy where they do not hold finished goods inventory unless there is foreseeable demand according to the warehouse's replenishment plan. There is a lead time to produce finished goods and the manufacturer will start pulling in raw material according to delivery requirements and production lead time. For example, if the production lead time is 2 periods and the delivery has to be made at time T , the manufacturer has to start pulling raw material into the production process at time $T-2$. The frozen production period, raw material availability, and production capacity add complications to production scheduling. The frozen production period is the period in the production schedule where no change in the schedule can be made (e.g. because raw material has already been pulled in). For example, a production manager may assign that the first F periods in the production schedule to be frozen. Raw material availability and production capacity limit the amount of output from production for any particular period. If raw material is not sufficient, the production to compensate for the insufficient amount can only be postponed to later periods when raw material becomes available. But if the production capacity is reached, the insufficient amount can be moved forward to be produced in earlier periods where there is some capacity left and material is available, and not in the frozen period, otherwise production of the remaining amount will have to be postponed to later periods. Output from forward production will be kept in finished goods inventory waiting to be delivered to the

warehouse. Output from the production schedule and finished goods inventory may not always meet the replenishment plan due to a number of constraints described above. The manufacturer also adopts an (s, S) policy for its raw material inventory management, however there is neither a minimum order quantity nor fixed order increment for raw material delivery. The procedure within the manufacturer and its interaction with warehouse and material provider are shown Figure 5.5.

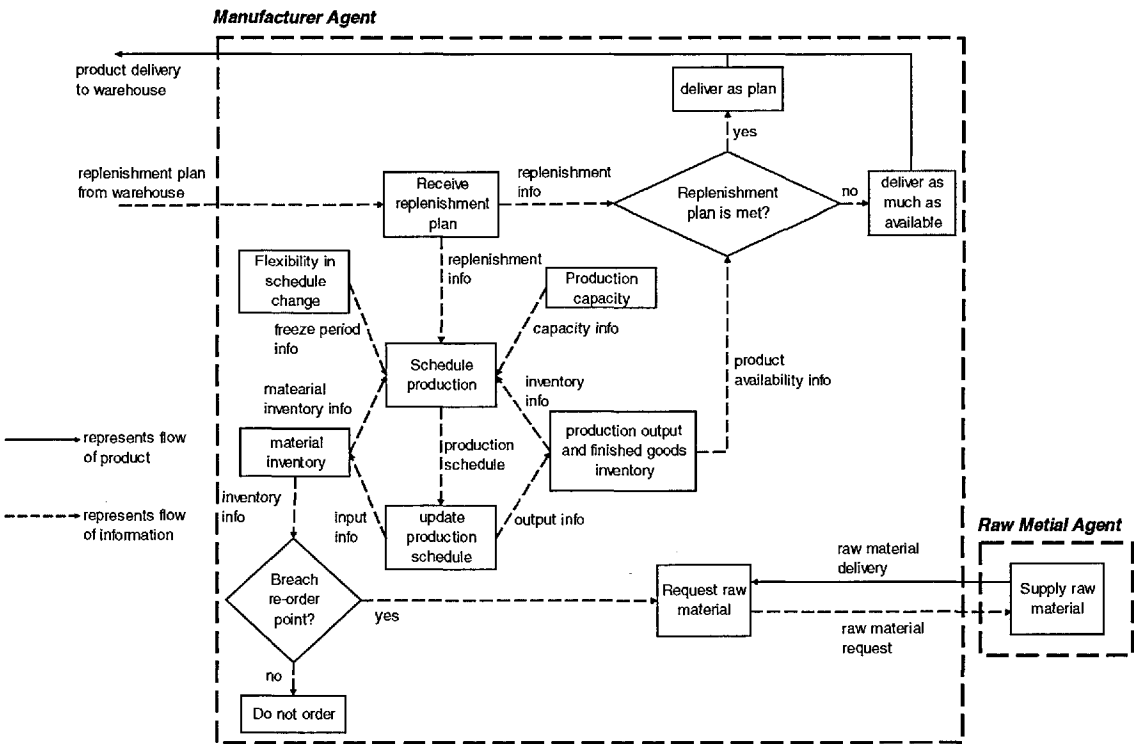


Figure 5.5: Manufacturer agent's supply chain policy and raw material agent

5.5 A Case study and Results

To implement our modelling concept, we consider a hypothetical case where there are 4 retailers, 1 warehouse, and 1 manufacturer in the supply chain system as previously shown in Figure 5.1. We are interested in how different fixed order increments (FOI) affect supply chain efficiency, complexity, and fill rate given the supply chain management policies in the previous section. We will also explore the relationships between different supply chain performance measures and transactional complexity for this case. We assume 1 period of simulation is equivalent to 1 day. Hence, the market will generate daily demand for each retailer, which in

our case study follows a uniform distribution of integers ranging from 0 to 69 (each with 1/70 probability). Settings for simulation of this case are shown in Table 5.2.

Agents	Parameters	Values
Market	Daily demand profile for each retailer	Uniform Integer [0, 69]
Retailer	Re-order point	200 units
	Order up to Level	400 units
	Minimum Order Quantity (RMOQ)	100 units
	Order Increment (RFOI)	<i>Varied</i>
Warehouse	Delivery lead time	3 days
	Base stock target	300 units
	Minimum Order Quantity (WMOQ)	100 units
	Order Increment (WFOI)	<i>Varied</i>
Manufacturer	Production capacity	250 units
	Production schedule frozen period	2 days
	Material re-order point	450 units
	Material order up to Level	900 units
Raw Material	Delivery lead time	2 days

Table 5.2: List of parameters setting for the case study

For each scenario according to the settings above, our simulation runs for 60 time periods. However, as the simulation outputs from early periods are usually highly dependent on the initial conditions (Law and Kelton, 2000), we disregard the outputs from the first 30 periods and calculate performance measures using the figures (e.g. inventory levels, number of orders, etc.) from the last 30 periods. We then simulate 20 repetitions for each scenario and report the average monthly performance. After simulating all scenarios for our case study, we summarise the results and examine the supply chain performance of the retailers and warehouse.

5.5.1 Retailers’ Performance

First we examine the retailers’ supply chain performance by varying the retailers’ order increment (RFOI) from 1 to 102 when the warehouse order increment (WFOI) is 1, 48, and 96, respectively. First we plot the RFOI against retailers’ entropic complexity, mean absolute percentage error (MAPE) of inventory, average number of orders, inventory turnover, and fill rate for each WFOI, as shown in Figure 5.6.

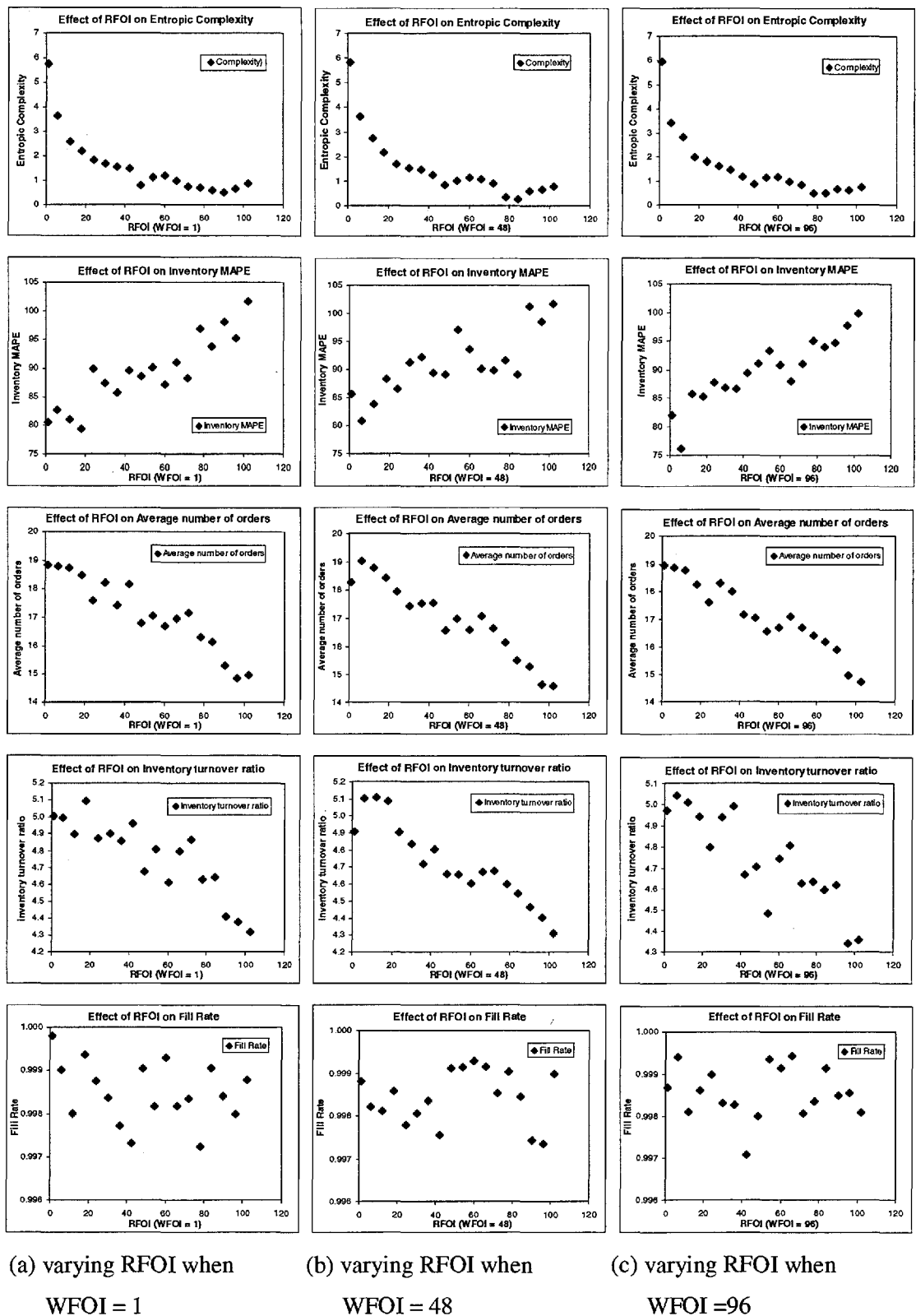
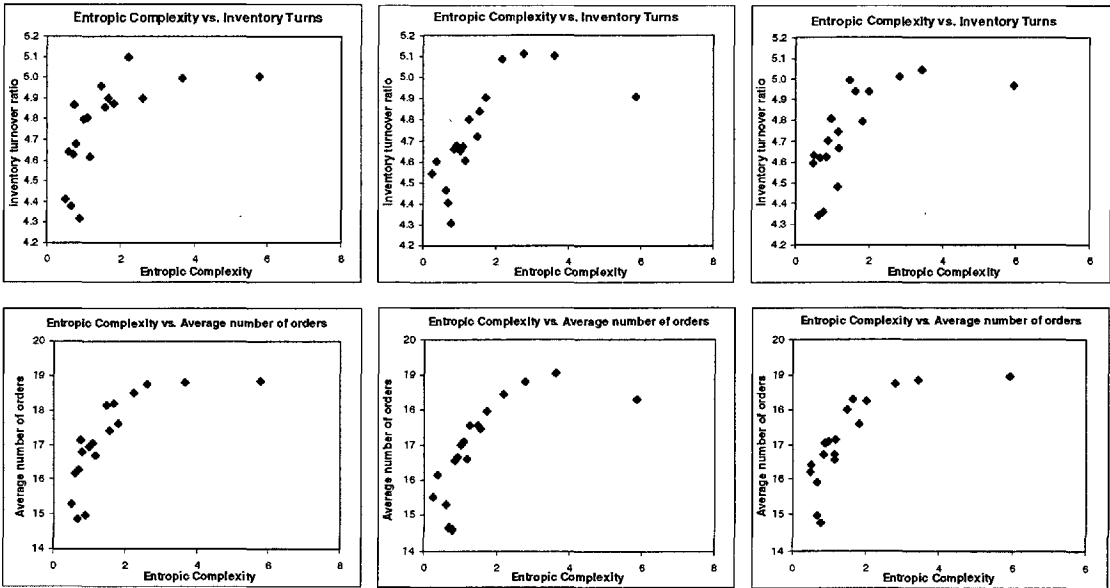


Figure 5.6: Effect of RFOI to retailers' performance

For MAPE of inventory which measures fluctuations in inventory level, we calculate this as the MAPE for inventory of the entire retailer echelon. From Figure 5.6, we can see that in this case

the retailers' fill rate (for customer demand) is not very sensitive to RFOI. We can also see that different choices of WFOI do not affect the pattern of relationships between RFOI and the retailers' supply chain performance. Entropic complexity which is a representative of complexity in the order making and order receiving system (e.g. in an IT system) decreases when RFOI increases. It first drops sharply but the rate of decrease declines as RFOI becomes larger. On the contrary, the MAPE which is representative of complexity in inventory management increases steadily and linearly as RFOI increases. This suggests that if we want to decrease order making/receiving complexity by increasing RFOI, we have to compromise on higher complexity in inventory management. The good news is that in our example RFOI does not have to be extremely high to reduce order making/receiving complexity since the entropic complexity decreases sharply when RFOI increases from a relatively small number. Both the inventory turnover ratio and average number of orders decrease quite linearly as RFOI increases. This again shows that increasing RFOI does not only increase inventory management complexity but also decrease inventory management efficiency; both are negative effects. On the other hand, increasing RFOI decreases order making/receiving complexity and increases order making/receiving efficiency (by reducing the average number of orders); both are positive effects. This can be a dilemma for a supply chain manager to find a balanced RFOI to obtain desirable supply chain performance both in inventory management and in order making/receiving transactions.

Next we examine the relationship among supply chain performance measures for the retailers. We plot inventory turnover ratio, average number of orders, and fill rate against entropic complexity and MAPE. The aim is to understand the potential benefit of more complex transactional systems. Figure 5.7 displays the results when WFOI is 1, 48, and 96.



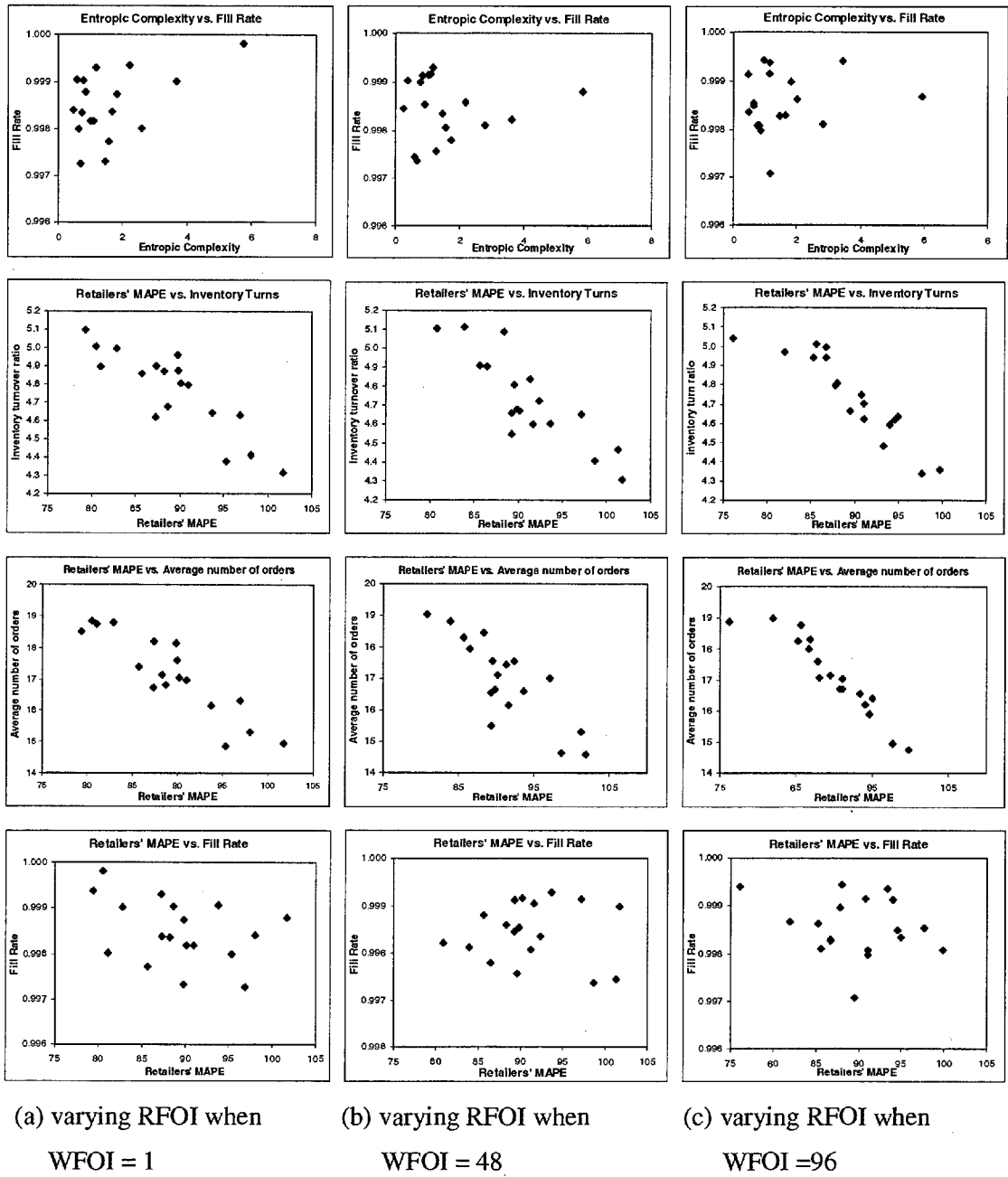


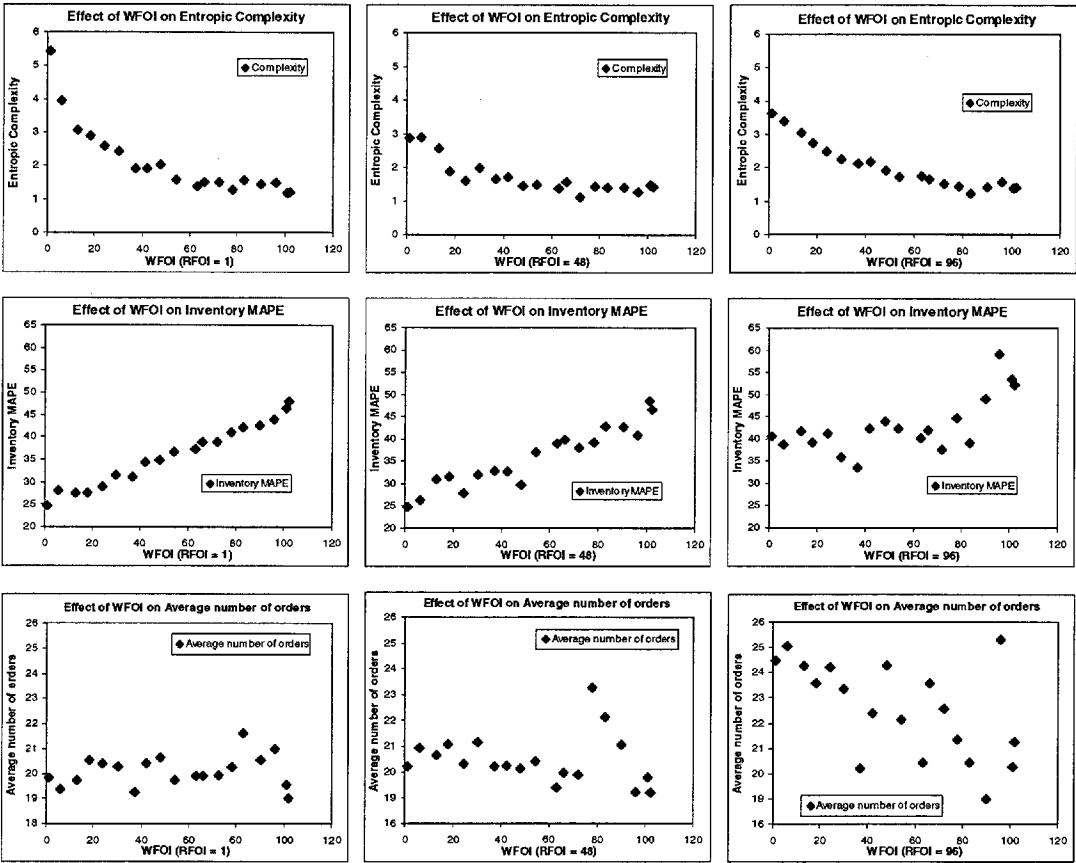
Figure 5.7: Relationships between retailers' supply chain performance measures

From Figure 5.7, no clear relationship could be observed regarding the fill rate. However, both inventory turnover ratio and average number of orders increase with entropic complexity, and decrease with MAPE. We can also see that the choice of WFOI does not affect the patterns of the relationships. These patterns are not a surprise given our previous observation about how RFOI affects supply chain performance. However, plotting these relationships can be useful because the patterns are quite strong. A supply chain manager can look at these plots or those relevant to his/her system to select an appropriate RFOI for his or her supply chain. For example, if he or she is interested in inventory turnover ratio and entropic complexity (e.g. they

are the most important performance measures for his or her supply chain), he or she can look at the plot *entropic complexity vs. inventory turnover ratio* and explore the trade-off between these two performance measures. By examining the plot the manager can see that he or she does not have to compromise a very high entropic complexity to improve inventory turnover ratio. The manager may choose to take a little more entropic complexity to improve inventory turnover and check with raw data to identify an appropriate RFOI. In other cases, the manager may learn from the plot that, from his or her current situation, accepting a higher entropic complexity improves inventory turnover only marginally (e.g. when the curve becomes flat).

5.5.2 Warehouse Performance

Next we examine the warehouse supply chain performance by varying the warehouse order increment (WFOI) from 1 to 102 when retailer order increment (RFOI) is 1, 48, and 96, respectively. We then plot the WFOI against warehouse's entropic complexity, mean absolute percentage error (MAPE) of inventory, average number of orders, inventory turnover, and fill rate for each RFOI, as shown in Figure 5.8.



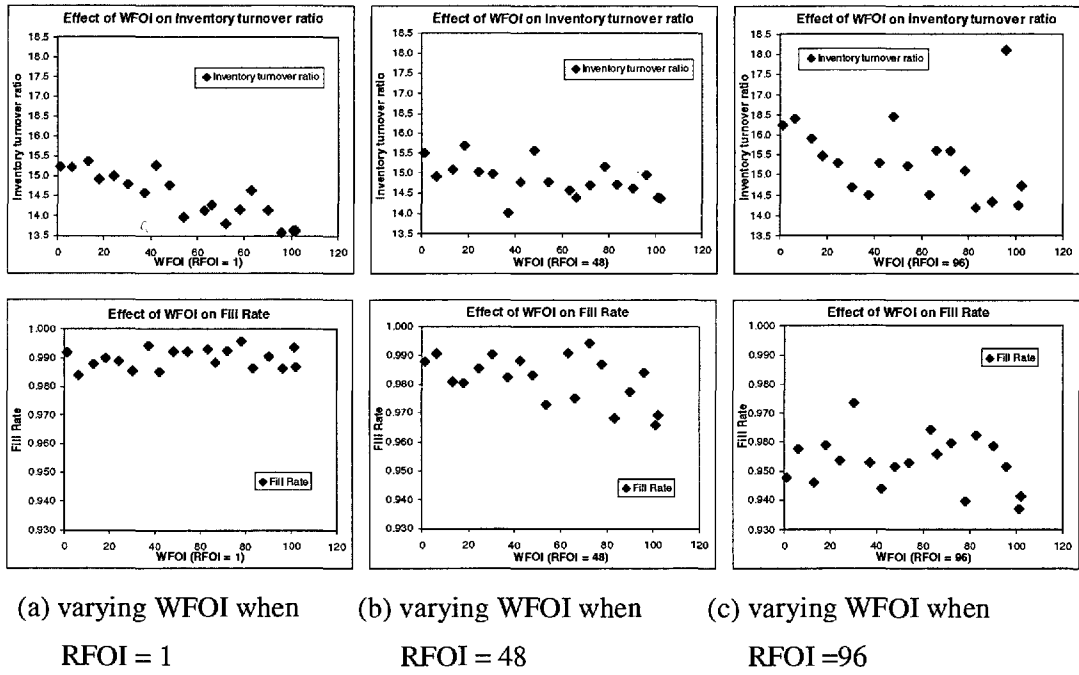
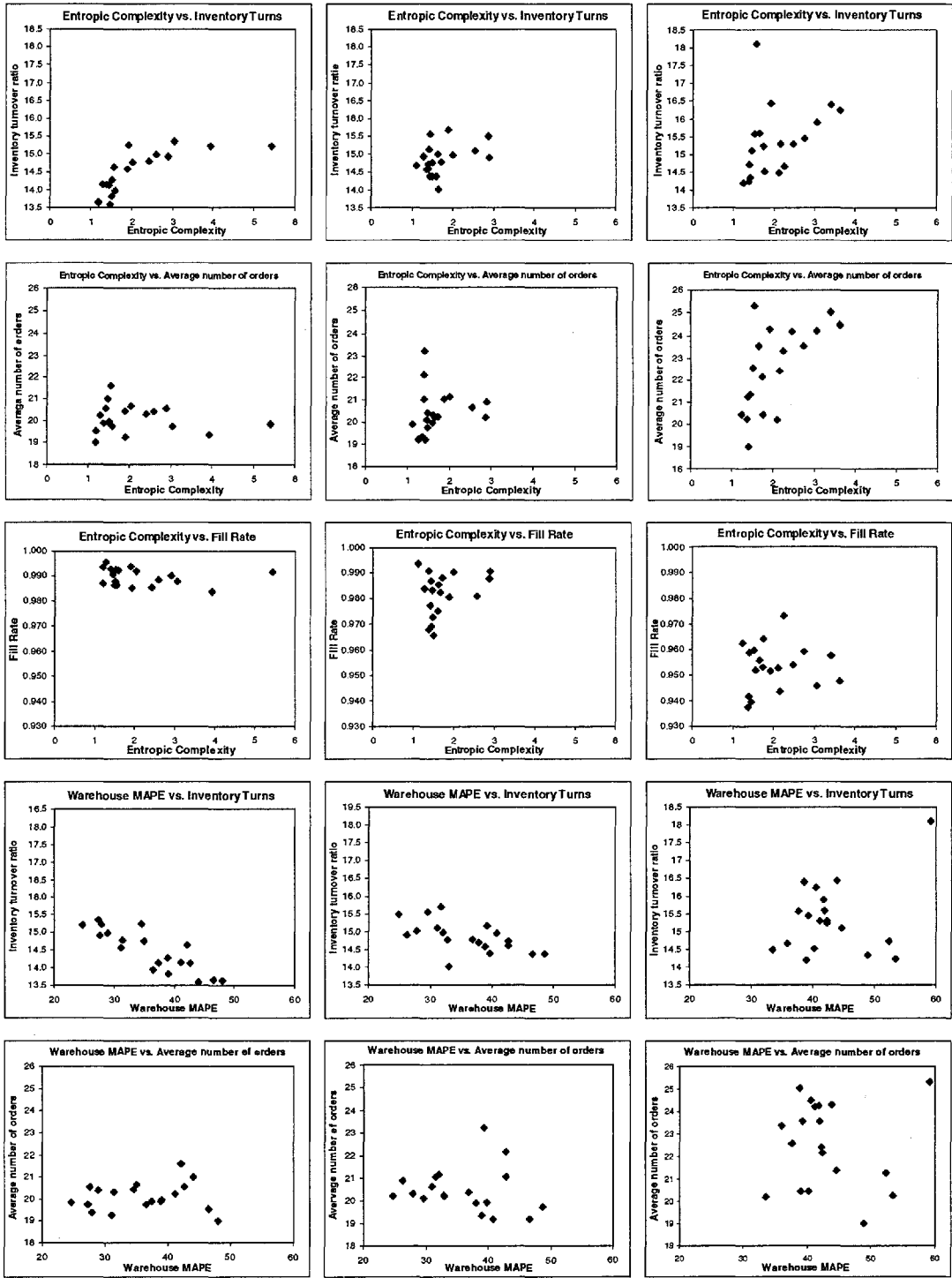


Figure 5.8: Effect of WFOI to warehouse performance

From Figure 5.8, we also see (similar to the retailers' case) that in this case the warehouse fill rate is not very sensitive to WFOI. However, in contrast, the warehouse fill rate is sensitive to RFOI. As RFOI becomes larger, the fill rate becomes worse. This is an interesting phenomenon revealing how the influence of constraints in the downstream supply chain propagates to the upstream supply chain. When RFOI = 1, the entropic complexity when WFOI is small is significantly greater than when RFOI is 48 or 96. However, when WFOI is large, the entropic complexity declines to about the same level along all choices of RFOI. Hence, the entropic complexity of the warehouse can be affected by both WFOI and RFOI. Again, this reveals propagation of the influence of downstream constraints to the upstream supply chain. However, the MAPE increases as WFOI increases and the trend is quite consistent along all choices of RFOI. We can see a slightly reducing trend in inventory turnover ratio as WFOI increases though the correlation is not very strong. Although the average number of orders decreases as WFOI increases when RFOI = 96, the trend when RFOI is 1 or 48 is not very clear. Therefore, the trend is not consistent along all choices of RFOI.

From Figure 5.8, it seems that the WFOI mainly affects complexity (entropic and MAPE) in the warehouse but its influence on warehouse efficiency, both in order making/receiving and in inventory, is not very strong. Hence, the selection of WFOI arguably depends more on the trade-off between order making/receiving complexity (entropic complexity) and inventory management complexity (MAPE).

Next we examine the relationships among supply chain performance measures in the warehouse. We plot inventory turnover ratio, average number of orders, and fill rate against entropic complexity and MAPE. Figure 5.9 displays the results when RFOI is 1, 48, and 96.



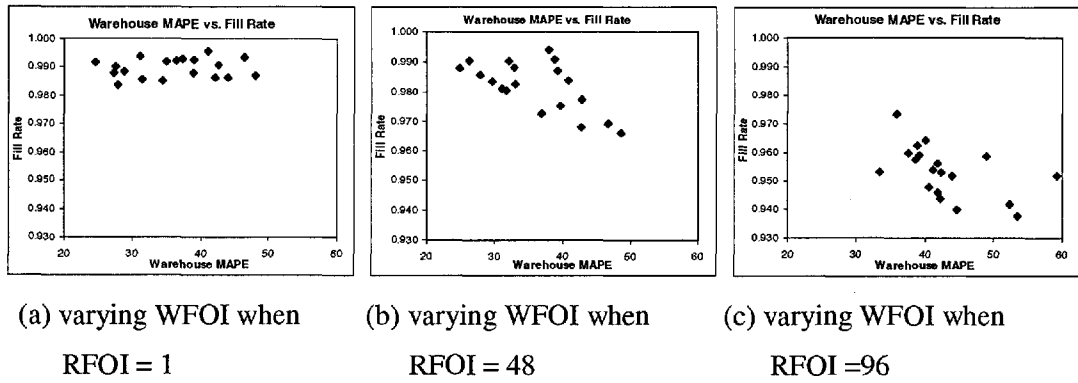


Figure 5.9: Relationships between warehouse supply chain performance measures

From Figure 5.9, similar to the retailers' case, no clear relationship between supply chain performance measures could be observed regarding the fill rate. However, among all choices of RFOI, the only noticeable relationship that we can observe is that of between entropic complexity and inventory turnover ratio. It is relatively consistent that inventory turnover ratio increases as entropic complexity increases, but the rate of increase in inventory turnover declines as entropic complexity becomes larger. Hence, there is no benefit in allowing very high levels of entropic complexity to improve inventory turnover at the warehouse. Other relationships are not very conclusive. Therefore, assumably the best course of action for a supply chain manager choosing WFOI is to select a range of acceptable inventory turnover and entropic complexity from the plot *entropic complexity vs. inventory turnover ratio* above and check with raw data to identify a set of feasible WFOIs. Then the manager can perform simulations based on these set of WFOIs to evaluate other performance measures before making any decision.

For both retailers and the warehouse, it is evident that fixed order increments strongly relate to inventory and transactional complexity in the supply chain. Although a higher level of complexity is normally undesirable, since other supply chain performance measures are interrelated with complexity, compromising on a higher level of complexity can be beneficial. Our results reveal the benefits, particularly regarding inventory turnover and average number of orders, gained from accepting higher level of complexity which in turn will generally result in more complex IT and transactional systems (e.g. for inventory management, purchasing, delivery). In this case study, we do an analysis only for the retailers' and warehouse's supply chain performance. However, if the effect of RFOI and WFOI on the manufacturer's performance is of interest, a similar analysis for the manufacturer can be performed.

5.6 Summary and Discussion

A supply chain is a discrete, non-linear, and complex system where each business entity or each group of entities are autonomous or semi-autonomous in decision making and policy setting. To assess supply chain operational performance measures such as efficiency, complexity, and fill rate, we propose a multi-agent based modelling approach to tackle the analysis of the system. It is widely documented that an IT revolution has taken place in this sector. We are interested in whether simulation can be used to assess the benefits of potentially more complex transactional systems. Due to its modularity, multi-agent based modelling can be easily modelled, implemented, and modified. However, its main strength is that the modelling framework is very well aligned with the structural and dynamic characteristics of supply chain systems. The results from our case study show that we can obtain several useful insights of supply chain behaviour from the outputs of multi-agent based simulation. In general, these behaviours are generally known as 'emergent behaviours' which are the outcomes of each agent's behaviour and their interaction with other agents in the supply chain. The results such as the trade-off between inventory management complexity (MAPE) and order making/receiving complexity (Entropic Complexity) and between inventory turnover ratio and average number of orders are very revealing and useful for supply chain managers. They are also powerful at convincing non-supply chain specialists of the validity of the results since they are the results of collective behaviours and interactions of agents in the supply chain, where the characteristics of each agent are relatively simple, visible, and easily understood. A supply chain manager can use these results to set a supply chain management policy to obtain desirable supply chain performance. A supply chain manager can also use these results as supporting material to negotiate with other business entities within the supply chain.

This modelling framework can be easily extended to assess efficiency, complexity, fill rate, and other supply chain performance measures in larger and more complex supply chain systems. For example, a modeller can include demand forecasting agents, multiple warehouses, multiple manufacturers, or extend the system further to incorporate upstream business entities. A modeller can also introduce into the model other uncertain elements in the supply chain system such as uncertainty in lead time, production yield, production process breakdowns, to name a few. In the next chapter, we will develop quantitative demand forecasting models based on historical sales data in order to improve demand forecasting accuracy and to obtain inputs for a distribution planning model in the chapter following.

Chapter 6

Demand Forecasting

This chapter discusses the case study at INEOS Fluor Ltd., a world leading fluorine-based products manufacturing company, which experiences highly fluctuating monthly product demand patterns throughout the year. In this chapter, we develop quantitative forecasting models based on INEOS Fluor's historical sales data and compare forecasting accuracy to current INEOS Fluor's current forecasting method. The results from demand forecasting can be used as inputs for a distribution planning model or other supply chain planning models. We also perform residual analysis to gain insights from distribution of residual errors. In this chapter, we apply existing quantitative forecasting techniques to a real-world industrial setting rather than attempt to develop a new forecasting methodology.

6.1 A case study with INEOS Fluor Ltd.

INEOS Fluor Ltd. is a world leading company focussing on the manufacture and supply of fluorine-based products, technologies, and services across a number of major industries, from pharmaceuticals to refrigeration, automotive to fluoropolymers. It is the world leader in refrigerants for car air-conditioning and propellants for medical inhalers. INEOS Fluor has 3 manufacturing sites in 3 continents to supply product demands globally. Due to the difference in manufacturing capability, certain products can only be produced at certain sites. Hence, to supply demand in one place of the products that can not be produced in that region, the products have to be shipped from one or both of the other two regions. INEOS Fluor also experiences highly fluctuating monthly product demand patterns throughout the year, especially for coolant related products. This is due to the fact that demand of these products is temperature dependent. The product demand will be at peak during summer and at trough during winter.

To plan production and materials purchasing, INEOS Fluor need to forecast their future product demands. This is normally done by the marketing team. The forecasting method adopted is a combination of sales force composite and customer surveys. INEOS Fluor has been in this business for a long time and has recorded historical monthly sales for several years. Given the availability of historical data, INEOS Fluor's management is interested in improving its forecasting capability through quantitative forecasting methods. After discussing with the Global Sales and Operations Manager and reviewing the historical data, we conclude that the time-series based models would be most appropriate as there is not enough information to construct an explanatory model. The historical database is also not large enough to construct the ARIMA model, hence we will focus on forecasting methods based on the Exponential Smoothing (ES) technique and the Time Series Decomposition (TSD) technique. To evaluate forecasting accuracy, we will use Mean Absolute Percentage Error (MAPE) in conjunction with Average Errors (to investigate forecasting bias).

6.1.1 Exponential smoothing technique

We will adopt three forecasting methods using an exponential smoothing technique. The mechanism of each method is explained below.

6.1.1.1 Simple Exponential Smoothing

The simple exponential smoothing technique is a popular forecasting method for stationary time series. A stationary time series is one which each observation fluctuates around a constant and the difference between the constant and the observation is considered as a random error with mean zero. The goal of this forecasting method is to distinguish between the random fluctuations and the basic underlying pattern. This is done by applying unequal weights to average (or to smoothen) the past data. The average will eliminate randomness, and relatively greater weights are given to the most recent data than the older data to reflect the assumption that the most recent data contain the most current information about what will happen in the future. The simple exponential smoothing forecasts can be expressed mathematically as follows:

$$F_t = \alpha X_{t-1} + (1 - \alpha)F_{t-1} \quad (6.1)$$

where $0 < \alpha \leq 1$ is the smoothing constant, which determines the relative weight placed on the most recent observation, F_t is the demand forecast for period t (New forecast), X_{t-1} is the actual observed demand of period $t-1$, and F_{t-1} is the forecast of demand made for period $t-1$ (Last

forecast). In general, F is the forecast and X is the actual observed demand for a period specified by the subscript.

Since $F_{t-1} = \alpha X_{t-2} + (1 - \alpha)F_{t-2}$, the formula can be rewritten as

$$F_t = \alpha X_{t-1} + \alpha(1 - \alpha)X_{t-2} + (1 - \alpha)^2 F_{t-2} \quad (6.2)$$

Finally, the formula can be expanded to show that the weights for older observations decay exponentially:

$$F_t = \alpha X_{t-1} + \alpha(1 - \alpha)X_{t-2} + \alpha(1 - \alpha)^2 X_{t-3} + \alpha(1 - \alpha)^3 X_{t-4} + \dots \quad (6.3)$$

The larger the α , the more weight is put on the most recent observation and the more sensitive the forecast to the latest change in data. To initialise the forecast, we start from $F_2 = X_1$ and the rest follow the formula above. A model user has to determine the α value. Here we will use an optimisation procedure to select α such that the MAPE is minimised.

6.1.1.2 Linear Exponential Smoothing (Holt's) Method

The linear exponential smoothing method or Holt's method accounts for time series with a linear trend. The method requires the specification of two smoothing constants, α for the value of the series and β for the trend (the slope). The equations are

$$S_t = \alpha X_t + (1 - \alpha)(S_{t-1} + G_{t-1}) \quad (6.4)$$

$$G_t = \beta(S_t - S_{t-1}) + (1 - \beta)G_{t-1} \quad (6.5)$$

where S_t is the value of the intercept at time t and G_t is the value of the slope at time t .

The τ -step-ahead forecast made in period t is then given as

$$F_{t,t+\tau} = S_t + \tau G_t \quad (6.6)$$

To initialise the process, we set $S_1 = X_1$ and $G_1 = X_2 - X_1$, and the rest follow the formula above. We also use an optimisation procedure to select α and β such that the MAPE is minimised.

6.1.1.3 Linear and Seasonal Exponential Smoothing (Winters's) Method

The linear and seasonal exponential smoothing method or Winters's method is capable of dealing with time series with both linear trend and seasonality. The forecast is based on three components which are the smoothed value of deseasonalised series (S_t), the smoothed value of trend (G_t), and the smoothed value of seasonal factor (c_t) as shown in the equations 6.7 – 6.9 below

$$S_t = \alpha(X_t/c_{t-H}) + (1 - \alpha)(S_{t-1} + G_{t-1}) \quad (6.7)$$

$$G_t = \beta(S_t - S_{t-1}) + (1 - \beta)G_{t-1} \quad (6.8)$$

$$c_t = \gamma(D_t/S_t) + (1 - \gamma)c_{t-H} \quad (6.9)$$

where H is the length of the season.

The τ -step-ahead forecast based on Winters's method made in period t can then be computed as

$$F_{t,t+\tau} = (S_t + \tau G_t)c_{t+\tau-H} \quad (6.10)$$

To initialise the forecast, we apply linear regression to the first $2H$ observations in the data series to obtain a regression model

$$y_t = a + bt \quad (6.11)$$

Next we set $G_i = b$, $S_i = a + G_i i$, and $c_i = \frac{X_i}{S_i}$, where $i = 1, 2, 3, \dots, H$. We then start the

forecast according to the Winters's method from period $H+1$ onward. We also use an optimisation procedure to select α , β , and γ such that the MAPE is minimised.

6.1.2 Time Series Decomposition technique

In time series decomposition, the pattern of the time series is broken down into Seasonality (S_t), Trend (T_t), Cycle (C_t), and Randomness (R_t). The forecast is then written as a function of these components as shown in equation 6.12.

$$F_t = f(S_t, T_t, C_t, R_t) \quad (6.12)$$

In general, the multiplicative model as shown in equation 6.13 is used to develop the forecast.

$$F_t = S_t \times T_t \times C_t \times R_t \quad (6.13)$$

Time series decomposition is especially appropriate for time series with seasonality. However, not all time series with seasonality exhibit trends or cycles. Below we propose three types of series and explain how to construct a forecasting model from them.

6.1.2.1 Decomposition Series with Seasonality only

This case assumes that there is no trend or cycle and that the forecast can be determined by a constant (average sales) and seasonal factors. Steps for decomposing this type of series are explained below.

1. Compute the sample mean ($\bar{X}$) of all the data
2. Divide each observation by the sample mean. This gives seasonal factors for each period of observed data.
3. Average the factors for like periods within each season. That is, average all the factors corresponding to the first period of a season, all the factors corresponding to the second period of a season, and so on. The resulting averages are the H seasonal factors.
4. Adding all seasonal factors together should be equal to H , otherwise we normalise the seasonal factors by multiplying them by H and dividing by the sum of original seasonal factors. The new seasonal factors will always add to exactly H . We can now forecast future sales as

$$F_t = S_t \times \bar{X} \quad (6.14)$$

where S_t is the seasonal factor for period t .

6.1.2.2 Decomposition of Series with Seasonality and Trend

For series with seasonality and trend, the trend component in the series is identified by H -period moving averages, where H is the length of the season. Below are steps for decomposing the series:

1. Calculate H -period moving averages starting from averaging sales of period 1 to period H , sales of period 2 to period $H+1$, and so on.

2. Assign each moving average to the centre period of the periods used to calculate the moving average. For example, if periods from 1 to 12 are used, the centre period is 6.5 (obtained by adding the index of the first and the last periods, then dividing by 2).
3. If the centre periods are not whole numbers, start from the first moving average, add a moving average to the adjacent one and divide by 2. Calculate a new centre period by adding the two original centre periods and divide by 2. Then assign the new moving average to the new centre period. For example, adding moving averages of centre periods 6.5 and 7.5 will result in a new moving average for a new centre period 7, adding moving averages of centre periods 7.5 and 8.5 will result in a new moving average for a new centre period 8, and so on.

$$MA_t = \frac{MA_{t-1/2} + MA_{t+1/2}}{2} \quad (6.15)$$

where t is a positive integer (e.g. 1, 2, 3, ...)

4. Divide each observation by the moving average of the corresponding period. Omit the observation where there is no corresponding moving average (the periods at the beginning and at the end of historical data). This gives seasonal factors for each period of observed data.

$$\frac{X_t}{MA_t} = \frac{T_t \times S_t \times R_t}{T_t} = S_t \times R_t \quad (6.16)$$

5. Average the factors for like periods within each season to eliminate the randomness. That is, average all the factors corresponding to the first period of a season, all the factors corresponding to the second period of a season, and so on. The resulting averages are the H seasonal factors.

$$\overline{S_t \times R_t} = S_t \quad (6.17)$$

6. Adding all seasonal factors together should be equal to H , otherwise we normalise the seasonal factors by multiplying each seasonal factor by H and dividing it by the sum of the original seasonal factors. The new seasonal factors will always add to exactly H .
7. Divide each observation by the seasonal factor for the corresponding period. This gives deseasonalised sales or the trend component (T_t) for each period. Assuming that the trend is linear, we can express the trend component as

$$T_t = a + bt \quad (6.18)$$

where a is a constant term and b is the slope of the trend

Using a linear regression technique to derive coefficients a and b such that the sum squared errors between the deseasonalised sales and the prediction by the regression model is minimised (see Appendix B), we can now forecast future sales as

$$F_t = T_t \times S_t = (a + bt) \times S_t \quad (6.19)$$

6.1.2.3 Decomposition of Series with Seasonality, Trend, and Cycle

In addition to trend and seasonal factors, the cyclical factors can also be decomposed from the data series. The deseasonalised moving averages include both Cycle and Trend. To identify the cyclical factors, we follow all the 7 steps in previous subsection and perform additional steps (Makridakis and Wheelwright, 1989) as explained below.

8. As shown in equation 6.20, deseasonalised sales consist of both Trend and Cycle components with random error.

$$\frac{X_t}{S_t} = \frac{T_t \times S_t \times C_t \times R_t}{S_t} = T_t \times C_t \times R_t \quad (6.20)$$

To eliminate the randomness from deseasonalised sales, a 3 x 3 moving average is used. This is done by taking first a three-period moving average of the deseasonalised sales and then another three-period moving average of these three-month moving averages. For example, first average the deseasonalised sales of periods 1, 2, and 3, and assign the moving average to period 2 (centre period), then do the same for periods 2, 3, and 4, and assign the moving average to period 3, and so on. Next, average the moving average corresponding to periods 2, 3, and 4, and assign it to period 3, and do the same for the rest. This will significantly reduce randomness in deseasonalised sales as shown in equation 6.21.

$$\overline{\overline{T_t \times C_t \times R_t}} = T_t \times C_t \quad (6.21)$$

However, there will be two 3 x 3 moving averages missing at the last two periods of the series. These values have to be estimated. For the value at period $N-1$, we will use a simple three-period moving average (average of deseasonalised sales in period $N-2$, $N-1$, and N).

$$MA_{N-1} = \frac{DS_{N-2} + DS_{N-1} + DS_N}{3} \quad (6.22)$$

For the value at the last period, there are two steps to obtain the estimate. First, calculate a simple two-period moving average using deseasonalised sales of the last two periods. This value will correspond to period $N-0.5$ (mid point of $N-1$ and N). Next, divide the latest change in moving average (moving average for period $N-1$ minus moving average for period $N-2$) by 2, and add it to the two-period moving average calculated earlier. This gives the estimate of the moving average for period N as shown in equation 6.23.

$$MA_N = \frac{DS_{N-1} + DS_N}{2} + \frac{MA_{N-1} - MA_{N-2}}{2} \quad (6.23)$$

9. Divide the moving average obtained in step 8 by the trend of the corresponding period from a regression model obtained in step 7. This gives the estimates of the cyclical factors as shown in equation 6.24.

$$\frac{MA_t}{a + bt} = \frac{T_t \times C_t}{a + bt} = \frac{T_t \times C_t}{T_t} = C_t \quad (6.24)$$

The forecasts can then be obtained according to equation 6.25.

$$F_t = (a + bt) \times S_t \times C_t \quad (6.25)$$

6.2 Forecasting Results

To perform the forecast using the forecasting methods mentioned above, an extensive amount of data processing and computation is required. To facilitate this computation, we created a template in a Microsoft Excel spreadsheet and code the computational procedure in *Visual Basic*. The advantages of a spreadsheet interface are that it is easy to use for people from both technical and non-technical backgrounds, and that most of the historical sales data are already stored in Microsoft Excel so the data transfer is convenient. The Microsoft Excel built-in *Solver* is used to perform optimisation for parameters in smoothing methods. The optimisation algorithm used is the Generalised Reduced Gradient (GRG2) nonlinear optimisation code developed by Lasdon *et al.* (1978). We tried this procedure with a different combination of starting points for parameters in smoothing methods and chose the best results (the lowest MAPE) for each method. We perform forecasting on one of the product at INEOS Fluor using their historical monthly sales data. To maintain business confidentiality, we call this product

“product A”, number each period instead of labelling actual month-year, and omit the number of sales volume. The results are displayed below.

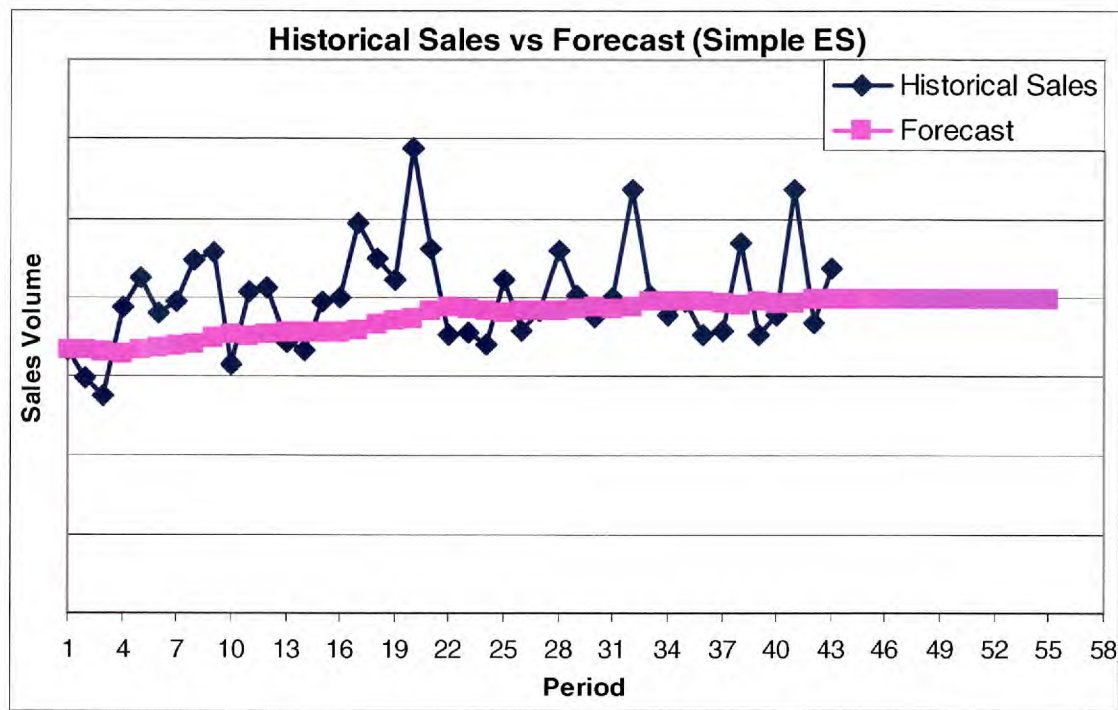


Figure 6.1: Forecasting results from Simple Exponential Smoothing
(Optimised $\alpha = 0.049567$)

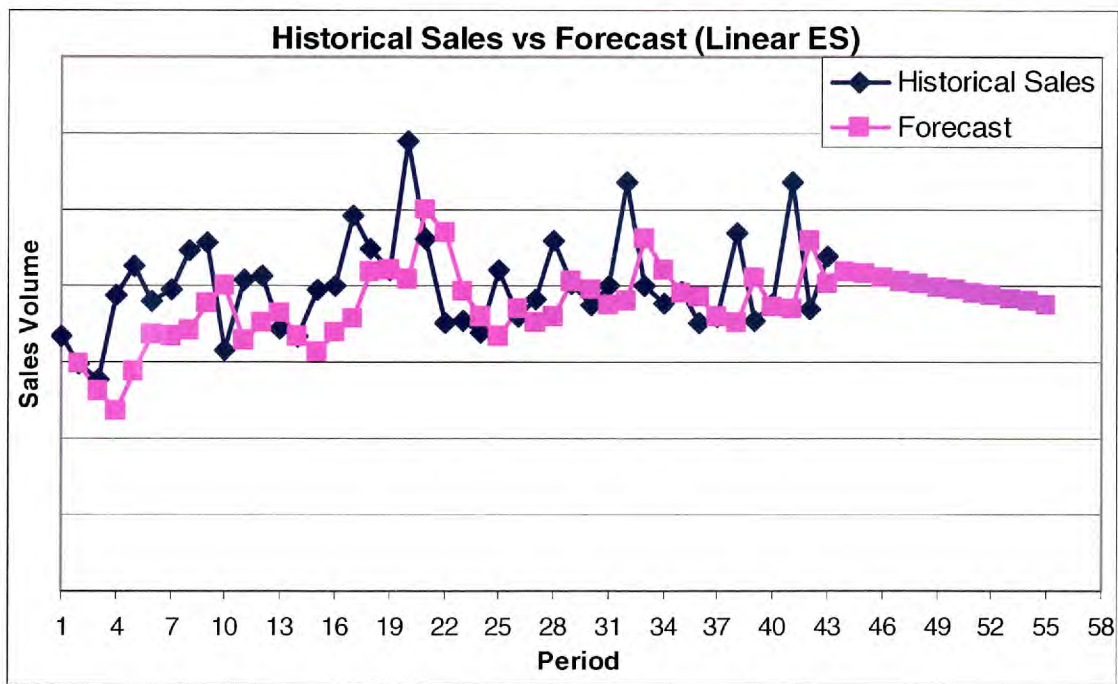


Figure 6.2: Forecasting results from Linear Exponential Smoothing
(Optimised $\alpha = 0.553275$, $\beta = 0.042921$)

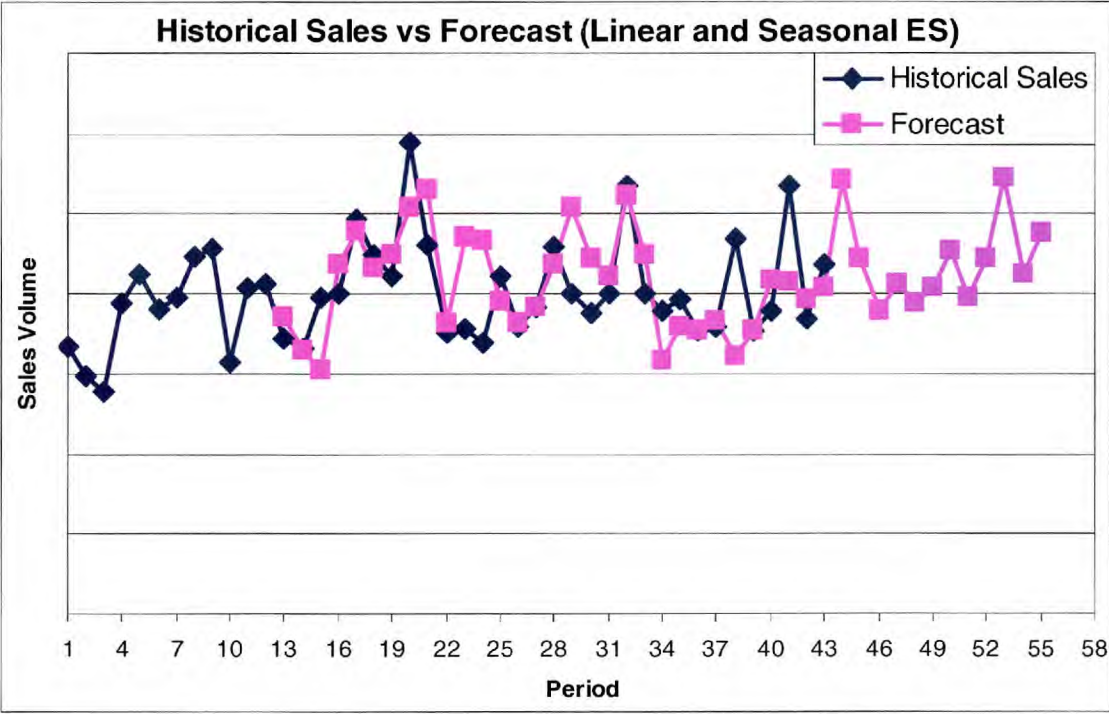


Figure 6.3: Forecasting results from Linear and Seasonal Exponential Smoothing
(Optimised $\alpha = 0.064755$, $\beta = 0.452319$, $\gamma = 0.572038$)

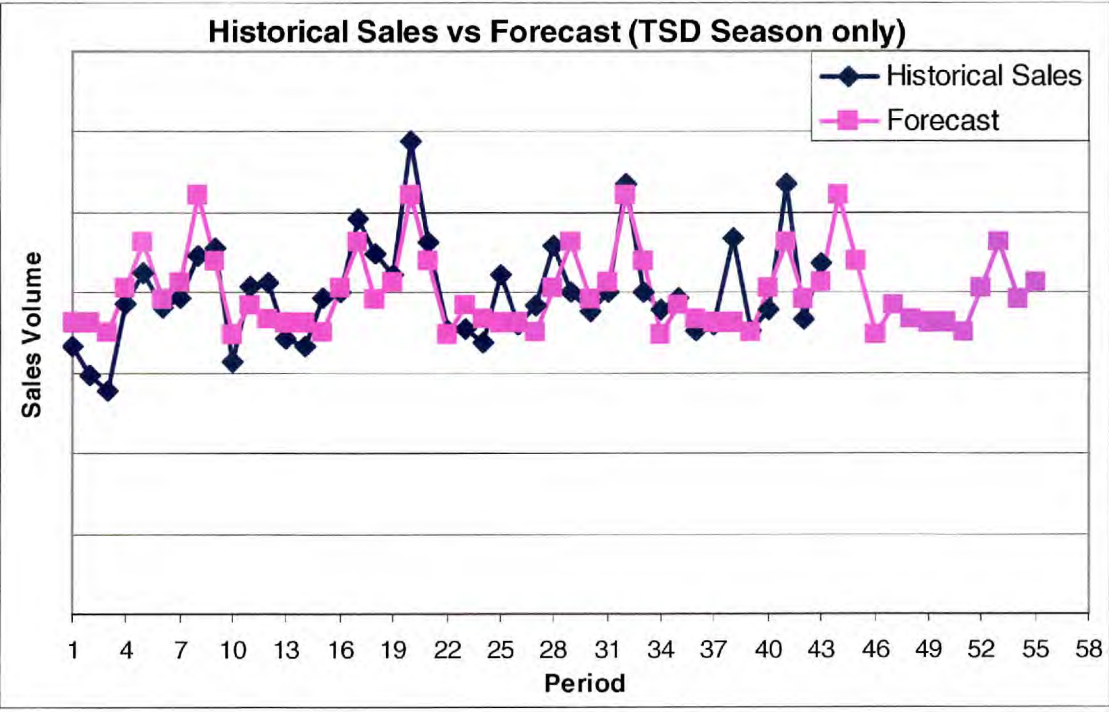


Figure 6.4: Forecasting results from TSD (Season only)

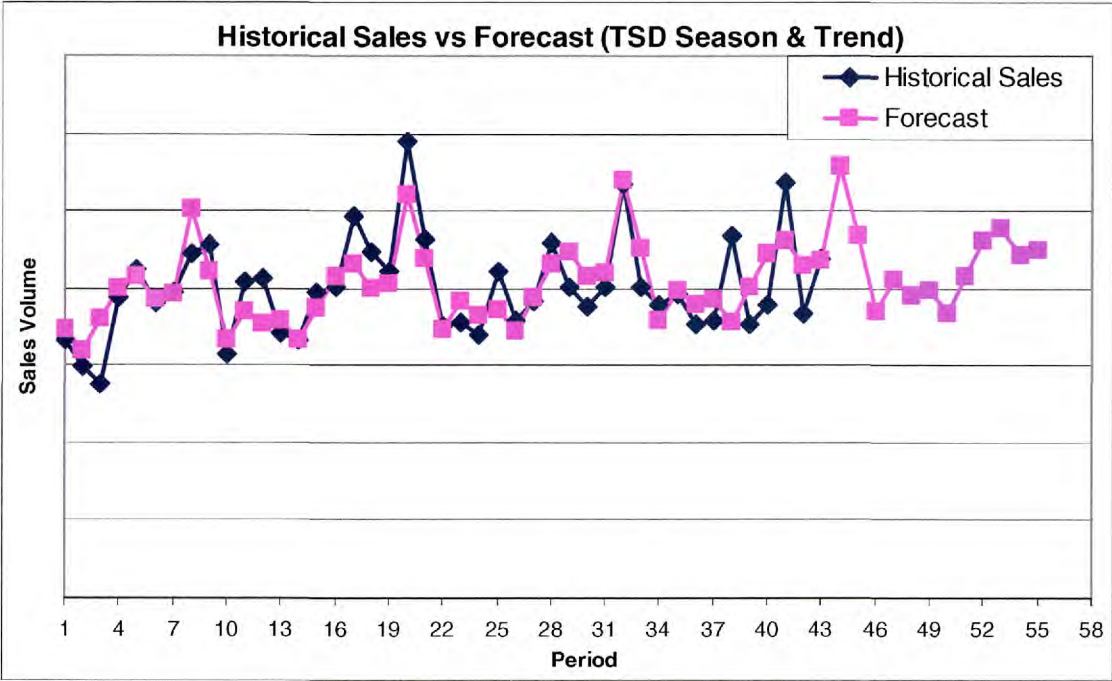


Figure 6.5: Forecasting results from TSD (Season and Trend)

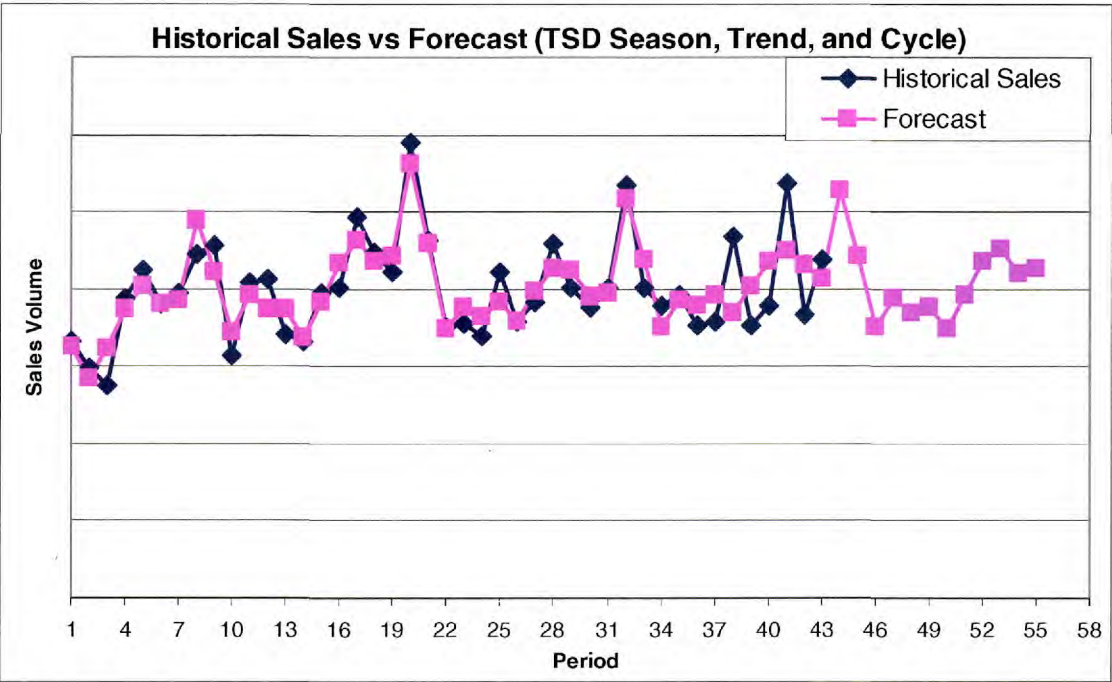


Figure 6.6: Forecasting results from TSD (Season, Trend, and Cycle)

The cyclical pattern of the product A time series is displayed in Figure 6.7.

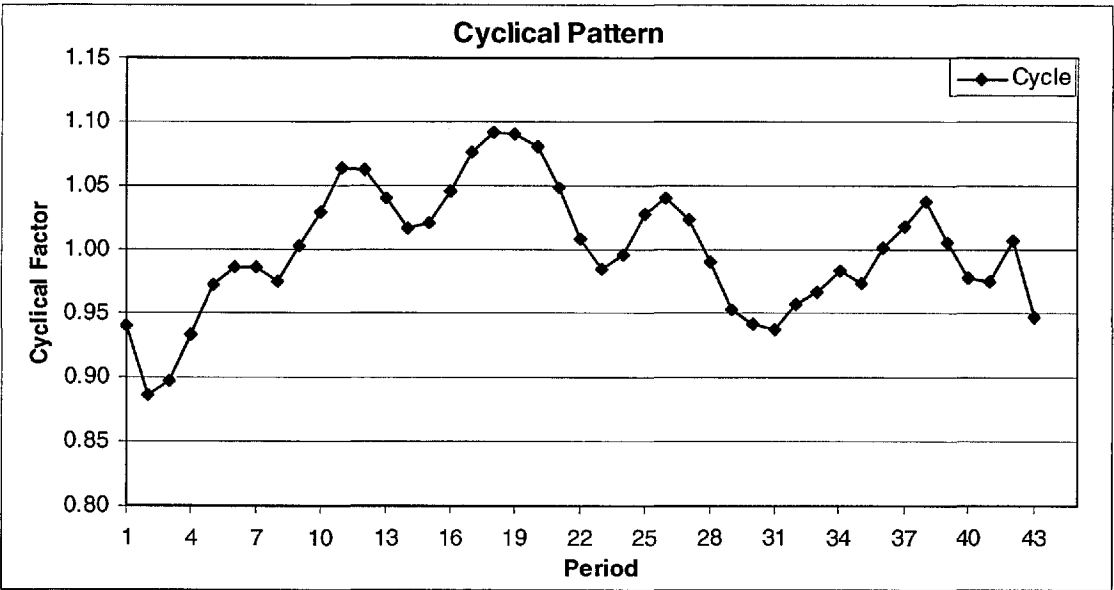


Figure 6.7: Cyclical pattern from TSD (Season, Trend, and Cycle)

A summary of forecasting accuracy of each method is shown in Table 6.1

Forecasting Method	MAPE	Avg Error	STD Error
ES (Simple)	12.98%	-16.10	132.71
ES (Linear)	14.88%	-64.70	149.99
ES (Linear and Seasonal)	10.96%	5.14	124.86
TSD (Season Only)	8.15%	-2.59	81.30
TSD (Trend & Season)	8.03%	0.36	82.99
TSD (Trend, Season, Cycle)	6.70%	0.17	68.63

Table 6.1: Summary of forecasting accuracy

The Mean Absolute Percentage Errors (MAPEs) of each method are compared as a bar chart in Figure 6.8.

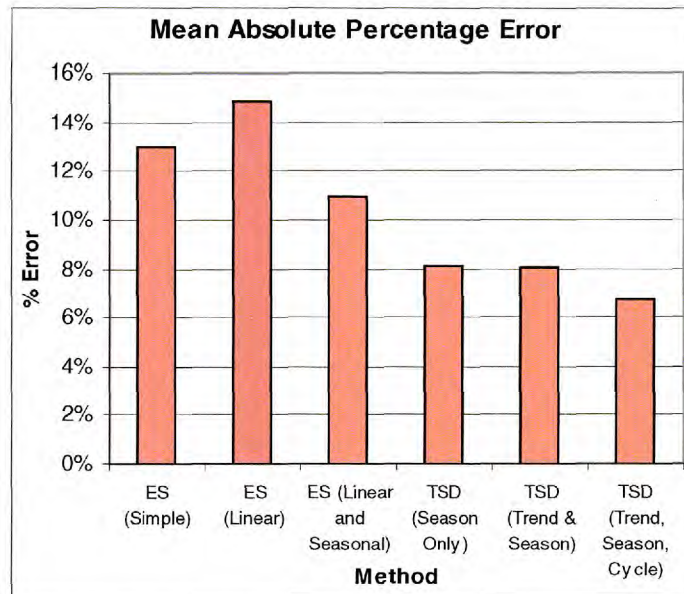


Figure 6.8: Mean Absolute Percentage Error bar chart

The average errors and standard deviation of errors of each method are compared as a bar chart in Figure 6.9.

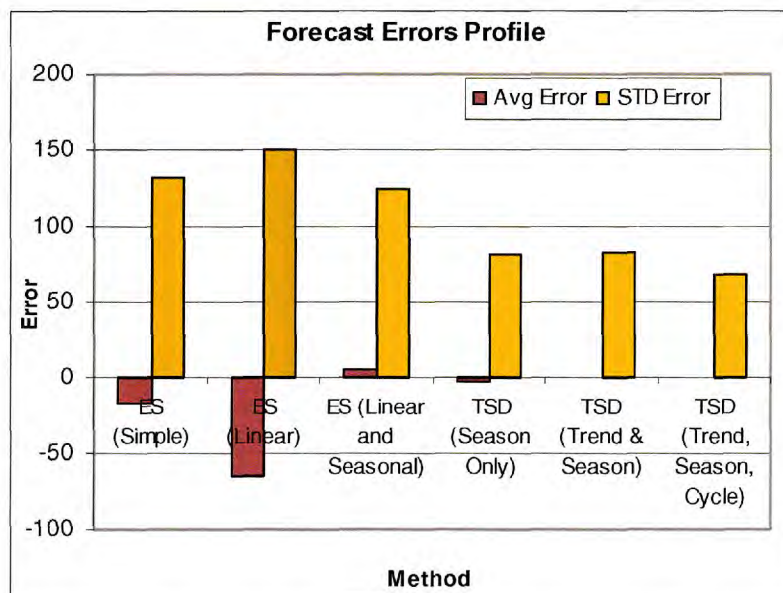


Figure 6.9: Average errors and standard deviation of errors bar chart

6.2.1 Comparison with INEOS Fluor’s Original Forecasts

We do a comparison with INEOS Fluor’s original forecasts to examine how the quantitative forecasts perform. Table 6.2 compares Absolute Percentage Error for the product A demand forecast for 9 periods. The Time Series Decomposition method uses historical sales data during

the past 26 periods. From the table, we can see that with only 26 data points (for seasonal length of 12 months) the TSD (season only) method performs much better than INEOS’s original forecasts (by sales force composite and customer survey). The MAPE (against actual sales realised during the 9 periods that the forecasts are made) of the TSD method is only 9.56%, whilst that of INEOS’s original forecasts is as high as 24.03%.

Month	Absolute Percentage Error	
	INEOS's Original Forecasts	TSD (Season only) Forecast
1	19.22%	13.25%
2	30.41%	14.96%
3	17.21%	13.27%
4	22.22%	9.33%
5	19.45%	0.86%
6	33.30%	4.21%
7	23.15%	13.48%
8	28.25%	12.63%
9	23.09%	4.06%
MAPE	24.03%	9.56%

Table 6.2: Absolute Percentage Error of INEOS’s forecasts and TSD forecasts

Figure 6.10 shows a comparison of product A sales forecasts for one market (labelled as “Market X”) made by INEOS’s original method and by the TSD (season only) method. The historical data are from the past 43 monthly sales, and the plots beyond period 43 are the forecasts. From the plot, the TSD forecasts match the historical sales pattern quite well, whilst INEOS’s original forecasts fluctuate more than the historical pattern.

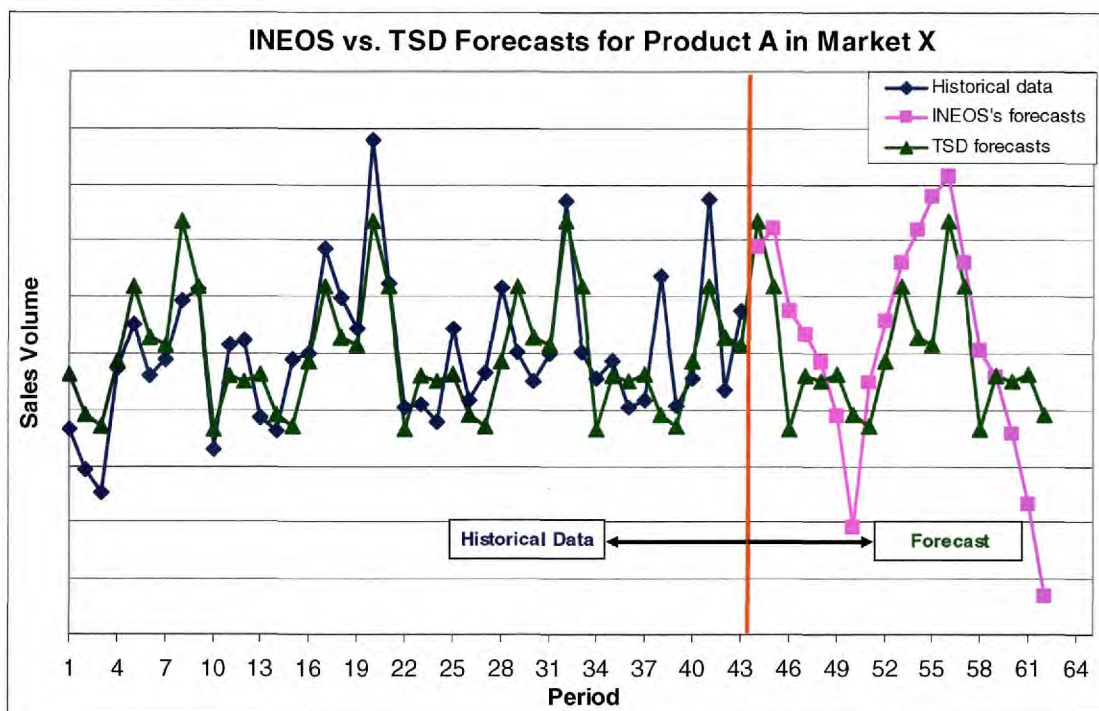


Figure 6.10: Product A sales forecasts for market X

Another example is Figure 6.11 which shows a comparison of product A sales forecasts for another market (labelled as “Market Y”) made by INEOS’s original method and by the TSD (season only) method. For this case, both the TSD forecasts and the INEOS’s original forecasts match the historical sales pattern well, however, INEOS’s original forecast is more optimistic for sales during the peak period.

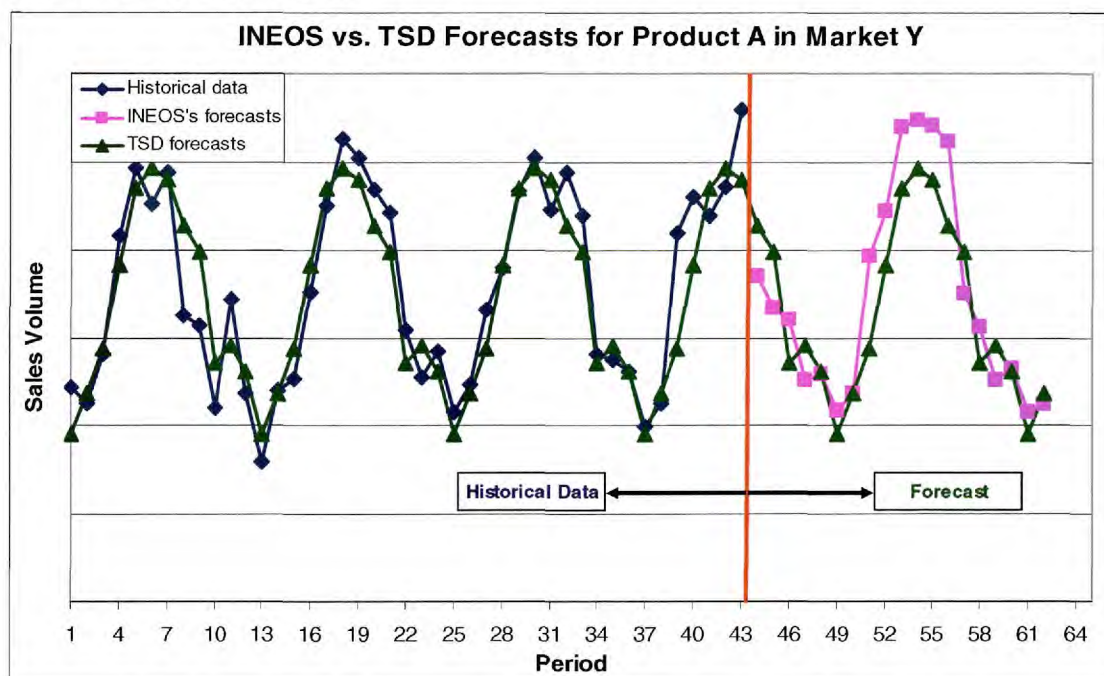


Figure 6.11: Product A sales forecasts for market Y

By plotting historical data, INEOS’s original forecasts, and TSD forecasts together, we can see how consistent the forecasts are with the patterns in historical data. Even if the quantitative methods have not yet been fully adopted, quantitative (TSD) forecasts can be used as a reality check for the original forecasts. If there is any significant difference between the two methods, the forecaster must check whether there is any specific reason why his or her forecasts should deviate from historical sales levels and patterns objectively identified by the quantitative forecasting methods.

6.3 Residuals Analysis

As discussed in section 6.3, it is desirable that the residual errors of the forecasting model be random and normally distributed with a mean of zero. The mean of residual errors is presented in Table 6.1. We can check the normal approximation of the residual errors by plotting histogram of the residual errors. In Figures 6.12 – 6.17, we present histograms of residual errors from each method. Normally distributed data should have a bell-shaped histogram.

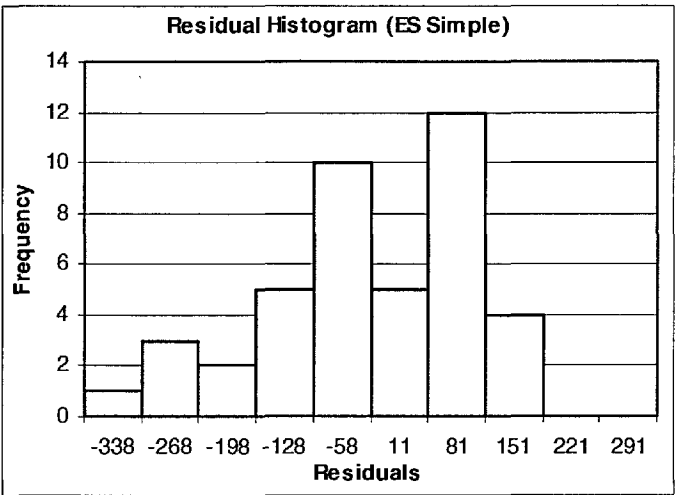


Figure 6.12: Histogram of residual errors from Simple Exponential Smoothing

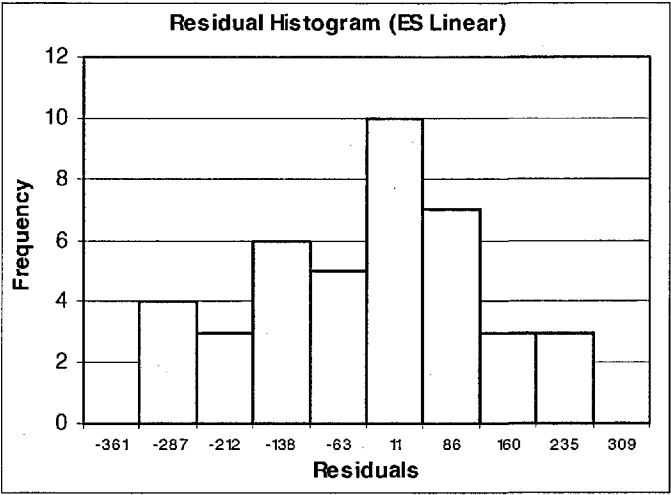


Figure 6.13: Histogram of residual errors from Linear Exponential Smoothing

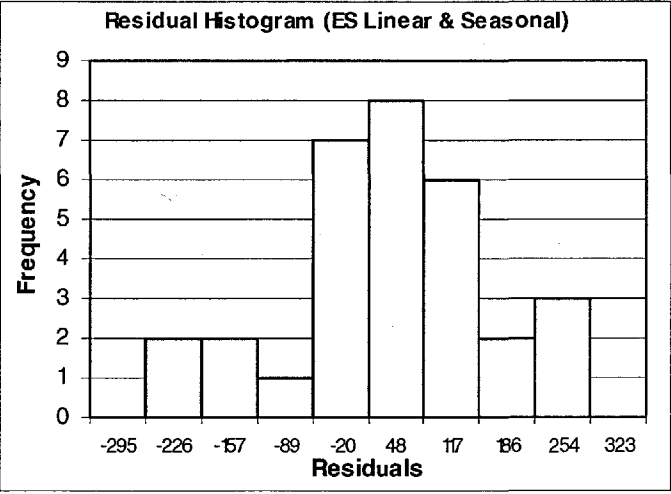


Figure 6.14: Histogram of residual errors from Linear and Seasonal Exp. Smoothing

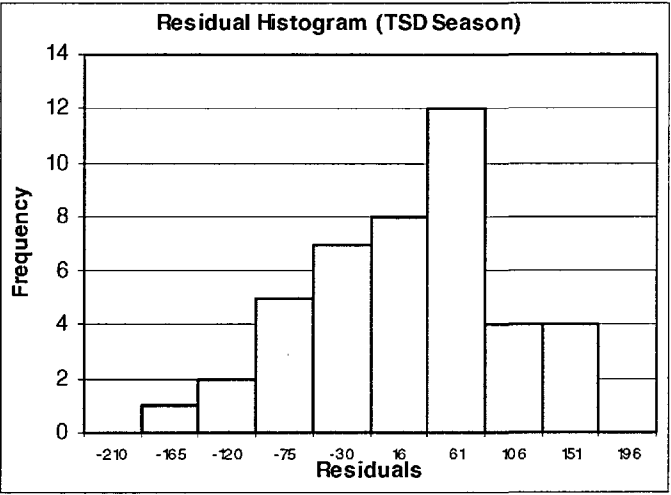


Figure 6.15: Histogram of residual errors from TSD (Season only)

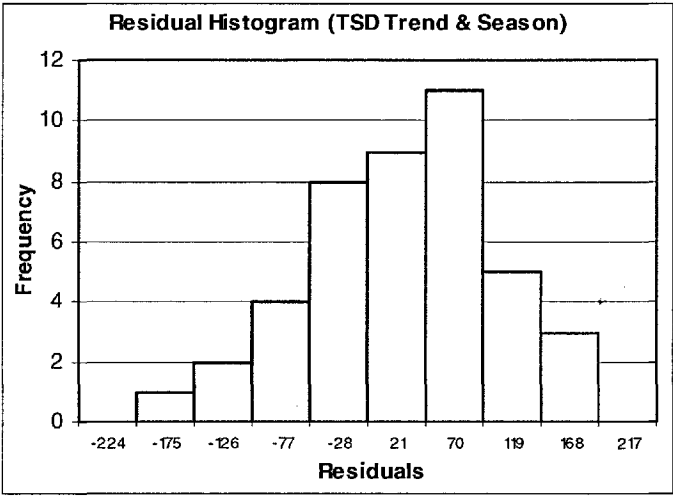


Figure 6.16: Histogram of residual errors from TSD (Trend and Season)

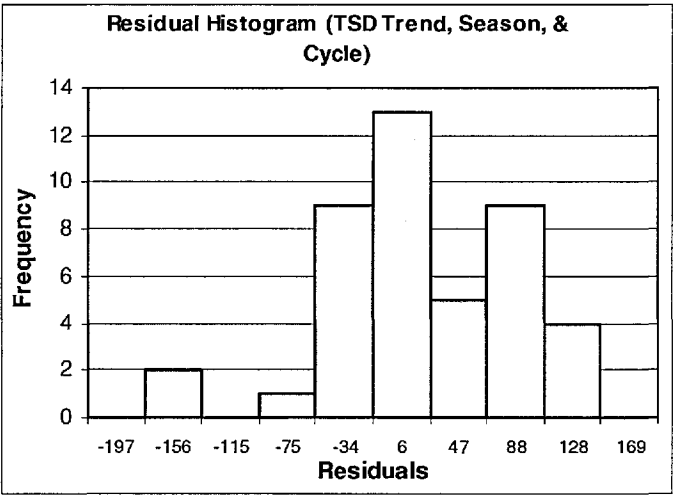


Figure 6.17: Histogram of residual errors from TSD (Trend, Season, and Cycle)

From all the histograms of residual errors, except a histogram in Figure 6.12 that skews to the left, all approximately have a bell shape. Therefore, we can conclude that there is no evidence suggesting that the residuals are not normally distributed. Generally, the shape of the histogram is good enough to infer a normal approximation. However, a more objective normality test, for example, using the Kolmogorov-Smirnov Test (see Levin and Rubin, 1997), can be performed if it is not clear from the shape of the histogram. The randomness of the residual errors can also be checked by calculating the autocorrelation of residual errors to make sure that there is no autocorrelation or by creating a run chart to make sure that there is no observable pattern in the residual errors. However, randomness of residual errors is neither strictly required for exponential smoothing methods nor for time series decomposition methods. Generally, having a distribution that appears normal with an average error close to zero is good enough.

6.4 Summary and Discussion

As almost all business operations require lead times, demand forecasting is extremely important in business planning. In this chapter, we review and present major qualitative and quantitative forecasting methods, characteristics of forecasts, and major forecasting accuracy indices. We perform a case study based on a real-world industrial setting at INEOS Fluor Ltd. where, traditionally, the demand forecasting is done by informal customer surveys and estimation based on experience by sales and marketing people. We introduced the INEOS Fluor management to the quantitative forecasting techniques: Exponential Smoothing technique and Time Series Decomposition technique, respectively. In comparison with the traditional practice, when appropriately selected, quantitative forecasting methods provide a significant improvement in forecasting accuracy. We also perform an analysis of residual errors from each forecasting method to check the distribution of the residual errors. For product A, only residual errors from simple exponential smoothing method exhibit non-normal-like distribution. It is also obvious from Figure 6.1 and the MAPE in Table 6.1 that the simple exponential smoothing method is not appropriate for demand forecasting of this product. Forecasters should use plots of historical sales vs. forecasts, MAPE, average residual errors, and residual errors distribution together to decide which forecasting method should be employed. For product A, all methods using Time Series Decomposition technique perform better than those based on Exponential Smoothing techniques. The performance among the three TSD techniques is not very different because the trend is not steep and the cyclical factors (Figure 6.7) vary in a relatively narrow range (0.9 – 1.1). Nevertheless, the TSD with Season, Trend, and Cycle method give the lowest MAPE. However, this is on the condition that a model user can correctly predict the cyclical factor at each period.

In the presence of lost sales, future demand can be underestimated when forecasts are based only on historical sales data (Nahmias, 1994). Therefore, if possible, historical *demands* (e.g. records of actual orders) rather than historical sales should be used. Finally, experience and other relevant information (e.g. a price reductions can lead to higher demand than usual) can also be used to adjust forecasting results from quantitative methods. In the next chapter, we will evaluate a product distribution strategy through a multi-agent based simulation model in which the results from our demand forecasting exercise will be used as the model input.

Chapter 7

Distribution Planning

Distribution is the delivery of products from one location to another. Distribution cost is one of the major costs in the supply chain. In some cases, it constitutes the majority of the supply chain cost. Generally, when designing a supply chain network, the distribution cost must always be taken into account since the supply chain network is the environment in which the distribution strategy is to be formulated (Chopra, 2003). Distribution planning decisions can be considered as both tactical level and operational level decisions (Simchi-Levi *et al.*, 2003) depending on which aspect of the distribution is being considered (e.g. daily vehicle routing is an operational decision, whilst deciding whether to engage a third party to handle delivery is a tactical decision).

Major considerations in distribution planning are distribution cost and service level. A supply chain manager has to maintain a high service level (on time delivery) at the lowest possible cost. These are often conflicting objectives. One way to reduce distribution cost is to adopt a “backhaul” strategy. Backhaul is the practice that containers or transporters of products sent from an origin to a destination unload the products at a destination point and load another product from the destination point to the origin point or other locations. However, managing distribution activities can also be a very complex task as a large amount of information has to be processed and several decisions have to be made in a short period of time. As an aid to a supply chain manager, in this chapter, we develop an agent-based simulation tool to reveal distribution costs (efficiency), service levels (robustness), and complexity of different distribution management alternatives. We will use a case at INEOS Fluor Ltd. to demonstrate the applicability of the simulation tool. Forecasted demands (such as those in chapter 6) will be used as one of the key inputs for the simulation model.

7.1 INEOS Fluor Ltd. Distribution Infrastructure

INEOS Fluor has manufacturing facilities in three major global markets. However, as INEOS Fluor's management do not want to disclose their distribution routes, we shall refer these three markets as regions X, Y, and Z, respectively. Each manufacturing site does not produce the same set of products. Some products with demand in region X but are not produced in region X may be produced in region Y or Z, and vice versa. For example, region X has a demand for product A, however product A is produced in regions Y and Z but not in region X. On the contrary, region Y has a demand for product B, however product B is produced in region X but not in region Y. Therefore, product A has to be transported from regions Y and Z to serve the demand in region X and product B has to be transported from region X to serve the demand in region Y. Figure 7.1 depicts the INEOS Fluor distribution infrastructure. The setting is from a real case, however, the product and region names are changed to maintain business confidentiality.

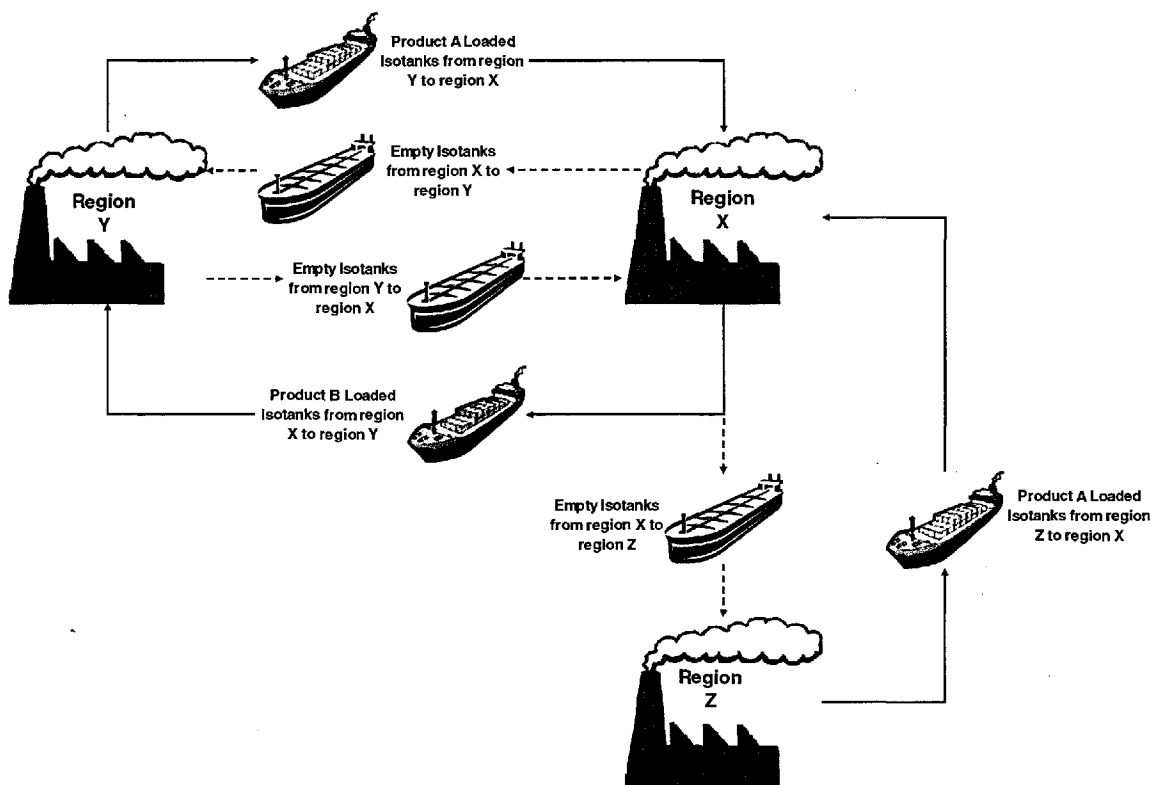


Figure 7.1: Diagram of product supply for product A and product B

The INEOS Fluor distribution cost includes the cost of leasing and operating the Isotanks, the bulk containers for chemical products. As INEOS Fluor's product sales are highly seasonal, the requirement for the Isotanks also varies considerably over the year. Since the Isotanks are sent regularly to and from regions X, Y, and Z, INEOS Fluor's management is interested in the

potential cost saving through backhauling. This means, for example, after a product is shipped from region X to region Y, the Isotanks that contain the product will not be returned empty but will be loaded with another product and sent back to region X or other destinations. By doing so, INEOS Fluor can save on the transportation cost of empty Isotanks. However, if there is no demand for products from region Y, the Isotanks will be kept there until product demand from region Y arises. A major drawback of the backhaul practice is that the Isotanks may have to stay and accumulate in one place for a long time and this can cause an Isotanks shortage in other locations. Therefore, the trade-off in this case is between the transportation cost saving and the service rate (on time delivery). We develop a multi-agent based simulation model to quantify and analyse this trade-off including the quantification of distribution system complexity.

7.2 A Simulation Model

According to the diagram in Figure 7.1, in the simulation model, there will be three locations for distributing and receiving the Isotanks. Each location represents both a manufacturing site and a market in region X, Y, or Z. The Isotanks will be circulated between these locations. The Isotanks are distributed weekly in fleets. However, the number of Isotanks in each fleet depends on the product demand and the number of Isotanks available in that period. The fleet size can be as small as one. Similar to chapter 5, the modelling approach is the multi-agent based simulation model. The modularity of the multi-agent based modelling approach is appealing because it is relatively easy to reconfigure the model if there is any change in the system. The *NetLogo* multi-agent based modelling platform (Wilensky, 1999), the same platform used in chapter 5, will be used. Figure 7.2 displays the diagram of the simulation mechanism.

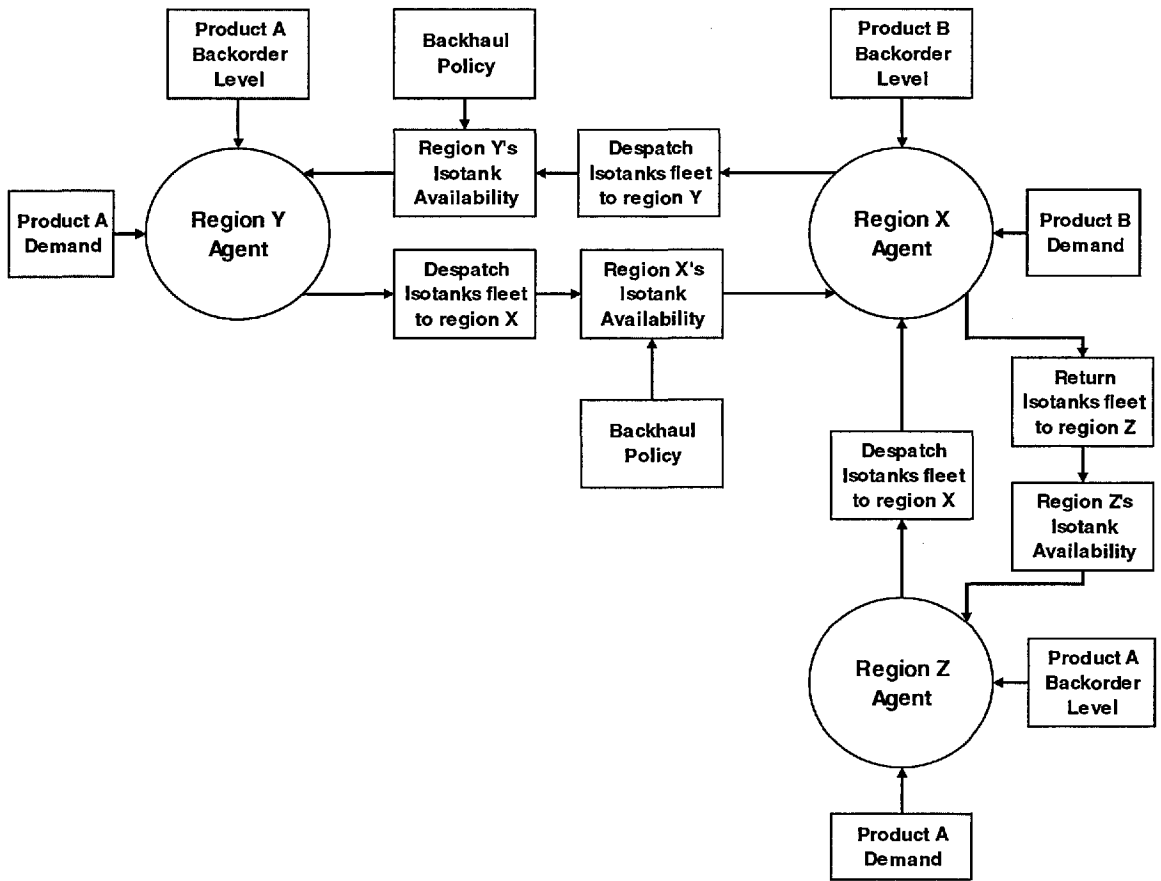


Figure 7.2: Diagram of the simulation mechanism

7.3 A case study

The simulation model is designed for the distribution of products A and B. Since region X's manufacturing site does not produce product A, the demand for product A in region X is served by the product A manufactured in regions Y and Z. We assume that the ratio of product A delivered for the European market by regions Y and Z is 2:1. In addition, the Isotanks containing product A from region Z will always be returned as empty Isotanks to region Z immediately after unloading. This means there is no backhaul on the X-Z route. On the contrary, the Isotanks containing product A from region Y can be backhauled by unloading product A at region X, reloading with product B, and heading back to region Y. The demand for product B in region Y is served by the product B manufactured in region X. The Isotanks containing product B from region X to region Y can also be backhauled. According to our case, Figure 7.2 depicts the user interface of the simulation model in *NetLogo*.

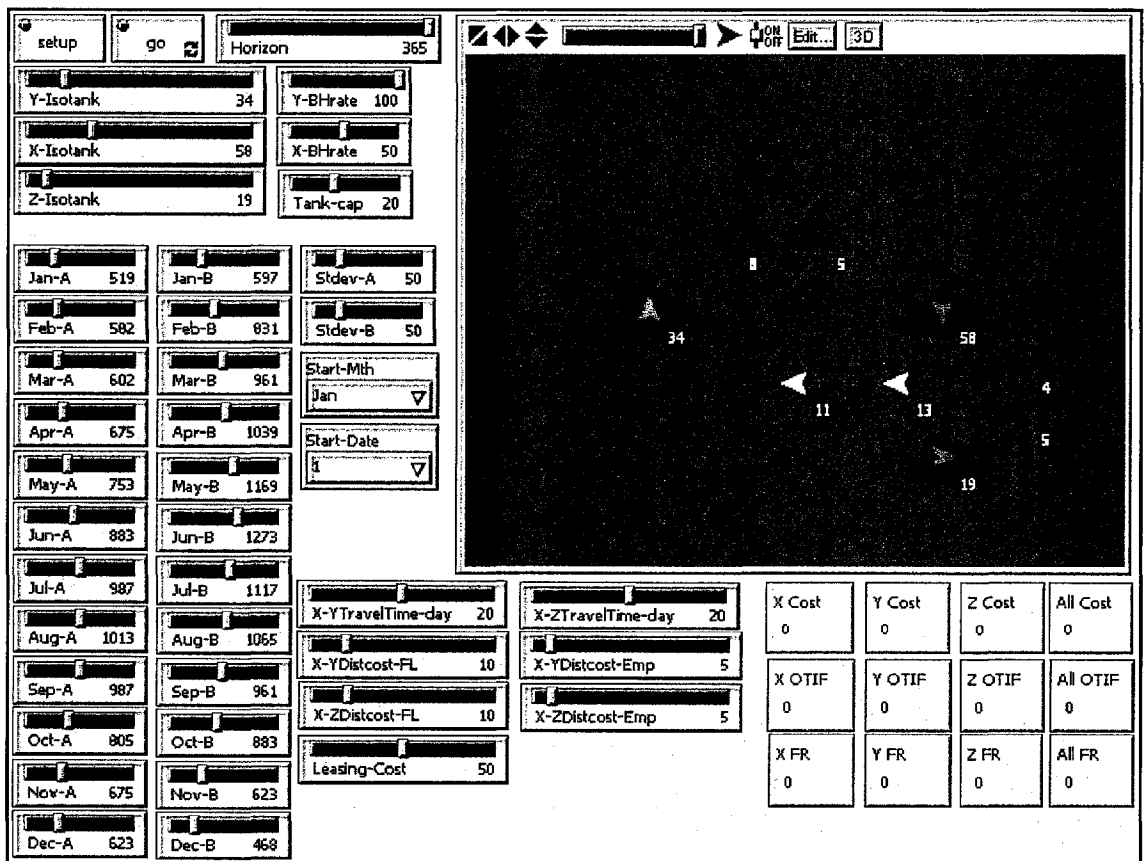


Figure 7.3: User interface on NetLogo

The simulation model is programmed to take the following inputs.

Simulation Horizon is the number of days to perform the simulation. The simulation will terminate after the last day of the simulation horizon.

Number of Isotanks at each location at the beginning of the simulation must be specified. These are the numbers of Isotanks available for use in regions X, Y, and Z.

Backhaul rate is the proportion of Isotanks that will be backhauled. There are two backhaul rates, one for the Isotanks shipped from region X to region Y and the other for the Isotanks shipped from region Y to region X. For example, if the backhaul rate for X-Y shipment is 70%, this means in a fleet of 10 Isotanks from region X to region Y, 7 of them will stay in region Y and be reloaded with product A before returning to region X, whilst the rest will be returned immediately to region X as empty Isotanks. The number of backhauled Isotanks for each fleet is rounded up or down to a whole number as appropriate (e.g. 70% of 8 Isotanks is 5.6 which will be rounded up to 6, 70% of 9 Isotanks is 6.3 and will be rounded down to 6). The fraction of the Isotanks being backhauled in each fleet may not be as accurate as specified,

however, in the long run the overall rate of Isotanks being backhauled will be close the specified figure.

Isotank capacity is the maximum quantity of the product that one Isotank can contain. If the Isotank capacity is set to 20 and the product demand in a particular week is 140, then 7 Isotanks will be needed. The number of Isotanks required will always be rounded up, for example, if the demand is 145, the number of Isotanks required is $145 \div 20 = 7.25 \approx 8$.

Monthly demand of product A in region X must be supplied. A model user can specify up to 12 months of demand. The time interval in which the demand is specified must correspond to the simulation horizon and the starting date of the simulation.

Monthly demand of product B in region Y must also be supplied and its time interval must agree with the simulation horizon and the starting date of the simulation.

Estimated standard deviation of monthly demand can be considered as an option to take into account the variability in demand when performing a simulation. A model user can set this to zero, if he or she does not want to add demand uncertainty. The standard deviation can be estimated from the standard deviation of forecasting errors as discussed in chapter 6.

Starting date and month of the simulation specify for which date and month the simulation shall start from. The product demands will then be generated with reference to the monthly demand profile and the day since the start of the simulation.

Isotanks travel time between locations determines the number of days required to transport Isotanks between regions X and Y, and between regions X and Z.

Distribution cost per one loaded Isotank is the cost of shipping a loaded Isotank between regions X and Y, and between regions X and Z.

Distribution cost per one empty Isotank is the cost of shipping an empty Isotank between regions X and Y, and between regions X and Z.

Since the product distribution in the simulation is performed weekly, the demand generated will also be the weekly demand. First, in each simulation day, we generate a daily demand Q , where Q is a random number that follows a normal distribution with mean

$$\mu_Q = \frac{D}{N}$$

and variance

$$\sigma_Q^2 = \frac{S^2}{N}$$

where D is the monthly demand of the month where that simulation day falls, S is the standard deviation of the monthly demand, and N is the number of days in the month where that day falls.

The weekly demand is then the cumulative daily demand in that particular week. This demand generating procedure applies to both product A and product B.

After running the simulation, the following outputs will be produced;

Distribution cost is the cost of transporting the Isotanks, both loaded and empty, from regions X, Y, and Z. In this model, the region of origin of the Isotanks being transported is the region that bears the distribution cost of those Isotanks. The overall distribution cost is the sum of all distribution costs.

Leasing cost is the cost of leasing cost of one Isotank per year. We assume that the leasing cost is proportional to the period of time in service of a particular Isotank.

On Time In Full (OTIF) rate is the rate (in percentage points) at which product orders are delivered on time and in full. The OTIF rate does not count any product which is delivered to partially fulfil the demand.

Fill rate is the fraction (in percentage points) of demand that is fulfilled on time. The fill rate counts both when the product is delivered on time and in full and when the product is delivered to partially fulfil the demand.

System Complexity level is the measure of complexity in the distribution network calculated according to the formula explained in section 7.4. This complexity is the overall system complexity looking at a global scale and is not subdivided into complexity attributed to each individual location.

In addition to the inputs via sliders and choosers on the user interface, the simulation model can also take inputs, via a text file, of Isotanks being in transit at the start of the simulation, Isotanks that will be available in the future during the simulation, and Isotanks that will be lost (e.g. end of leasing contract) in the future during the simulation. Figure 7.4 shows example of the format of the inputs via a text file.

```
"List of Z to X fleet size"
4 5

"List of Z to X arrival time (day)"
7 14
```

Figure 7.4: Format of inputs in a text file

In the text file, starting from the earliest arrival, available, or lost date, a user must list the number of Isotanks fleet size corresponding to the future arrival, available, or lost date (the date from the beginning of the simulation).

7.4 Costs, Service Levels, and Complexity Calculation

The distribution cost is the cost of transporting loaded Isotanks and empty Isotanks. Equation 7.1 shows the distribution cost for region X.

$$DC_X = \sum_i Ef_i Ec_{x-y} + \sum_j Lf_j Lc_{x-y} + \sum_k Ef_k Ec_{x-z} \quad (7.1)$$

where DC_X is the distribution cost for region X, Ef_i is the fleet size of empty Isotanks to region Y fleet i , Ec_{x-y} is the cost of distributing one empty Isotank between regions X and Y, Lf_j is the fleet size of loaded Isotanks to region Y fleet j , Lc_{x-y} is the cost of distributing one loaded Isotank between regions X and Y, Ef_k is the fleet size of empty Isotanks to region Z fleet k , and Ec_{x-z} is the cost of distributing one empty Isotank between regions X and Z.

Equation 7.2 shows the distribution cost for region Y.

$$DC_Y = \sum_p Ef_p Ec_{x-y} + \sum_q Lf_q Lc_{x-y} \quad (7.2)$$

where DC_Y is the distribution cost for region Y, Ef_p is the fleet size of empty Isotanks to region X fleet p , Ec_{x-y} is the cost of distributing one empty Isotank between regions X and Y, Lf_q is the

fleet size of loaded Isotanks to region X fleet q , Lc_{x-y} is the cost of distributing one loaded Isotank between regions X and Y.

Equation 7.3 shows the distribution cost for region Z.

$$DC_Z = \sum_r Lf_r Lc_{x-z} \quad (7.3)$$

where DC_Z is the distribution cost for region Z, Lf_r is the fleet size of loaded Isotanks to region X fleet r , Lc_{x-z} is the cost of distributing one loaded Isotank between regions X and Z.

The overall distribution cost is the cost of transporting the Isotanks between all locations plus the leasing cost of the Isotanks in the distribution system as shown in equation 7.4.

$$DC_{Overall} = DC_X + DC_Y + DC_Z + LC \quad (7.4)$$

LC is the leasing cost of all Isotanks in the distribution system and can be calculated as follows;

$$LC = \frac{Tc}{365} \sum_s Tp_s \quad (7.5)$$

where Tc is the leasing cost per year per Isotank, Tp_s is the duration (number of days) in which the Isotank s is used within the distribution system, and 365 is the number of days in one year.

There are two indices for service levels: OTIF rate and Fill rate. These indices measure the robustness of the distribution system as well as customer satisfaction. We assume that the on time delivery is subject only to the availability of the Isotanks when despatching the product (e.g. manufacturing output is always on time or finished goods inventory is always sufficient). If the Isotanks available are not sufficient for despatching the product to fill demand in full, all available Isotanks will be used and the remaining product demand is backlogged. Hence, the On Time In Full (OTIF) rate can be calculated as follows;

$$OTIF = \frac{\text{Number of orders delivered on time and in full}}{\text{Total number of orders}} \times 100 \quad (7.6)$$

Next, the Fill rate (FR) can be calculated as follows;

$$FR = \frac{\text{Quantity of products delivered on time (both in full and in part)}}{\text{Total demand of products}} \times 100 \quad (7.7)$$

Similar to vehicle routing complexity in chapter 4, we will use the entropic complexity as a measure of distribution system complexity. Entropic complexity represents the amount of information required to describe the system. A high complexity will generally negatively impact the performance as it may be more difficult to manage the system (Guimaraes *et al.*, 1999). The calculation of complexity is shown in equation 7.8.

$$H_{DS} = - \sum_i \sum_j \sum_k \sum_l p_{ijkl} \log_2 p_{ijkl} \quad (7.8)$$

where H_{DS} is the complexity level of the distribution system, p_{ijkl} is the probability that any Isotank fleet will be loading product i (product A, product B, or Empty), starting a journey from location j (X, Y, or Z), heading to location k (X, Y, or Z), and having a fleet size of l .

7.5 Scenarios Setting

To perform a case study, we use the setting as shown in Table 7.1 which approximately reflects the real situation at INEOS Fluor.

Input Parameters		Value
Simulation Horizon		365 days
Number of Isotanks available at the beginning of the simulation	Region X	58 tanks
	Region Y	34 tanks
	Region Z	19 tanks
Isotank Loading Capacity		20 SKUs
Start Month		January
Start Date		1
Traveltime	Between X and Y	20 days
	Between X and Z	20 days
Distribution cost per one loaded Isotank	Between X and Y	10 MUs
	Between X and Z	10 MUs
Distribution cost per one empty Isotank	Between X and Y	5 MUs
	Between X and Z	5 MUs
Leasing Cost per Isotank per year		50 MUs
Standard Deviation of Monthly Demand	Product A (for region X's market)	50 SKUs
	Product B (for region Y's market)	50 SKUs

Table 7.1: Parameters input for a case study simulation

Here, SKU means Stock Keeping Unit and MU means monetary unit.

Table 7.2 shows the monthly demand for product A in region X and product B for region Y. The monthly demand can be obtained from forecasting results such as those in Chapter 6.

Month	Product A Demand (for region X's market)	Product B Demand (for region Y's market)
January	519	597
February	582	831
March	602	961
April	675	1039
May	753	1169
June	883	1273
July	987	1117
August	1013	1065
September	987	961
October	805	883
November	675	623
December	623	468

Table 7.2: Monthly demands for product A and product B

In addition, we also specify the Isotanks in transit at the beginning of the simulation, the number of Isotanks that will be available in the future during the simulation, and the number of Isotanks that will be lost in the future during the simulation in Tables 7.3, 7.4, and 7.5.

Origin	Destination	Product	Fleet Size	Arrival Date
Z	X	Product A	4	7
Z	X	Product A	5	14
Y	X	Product A	5	9
Y	X	Product A	8	15
X	Y	Product B	11	10
X	Y	Product B	13	17

Table 7.3: Details of Isotanks in transit at the beginning of the simulation

Location	Number of Isotanks	Available Date
X	5	5
X	2	17
Y	2	5
Y	3	17
Z	2	5
Z	1	17

Table 7.4: Details of Isotanks that will be available in future date at each location

Location	Number of Isotanks	Lost Date
X	2	3
X	3	12
Y	2	3
Y	3	12
Z	2	3
Z	3	12

Table7.5: Details of Isotanks that will be lost in future date at each location

The decision being considered by INEOS Fluor is whether to adopt the backhaul strategy. To simulate the outcomes of different alternatives, we simulate 5 different scenarios of backhaul rates as shown in Table 7.6.

Scenario	Region X's Backhaul rate	Region Y's Backhaul rate
A	0%	0%
B	50%	50%
C	100%	50%
D	50%	100%
E	100%	100%

Table 7.6: Details of five different backhaul strategies

The region X's Backhaul rate is the backhaul rate of Isotanks fleet despatched from region X. The region Y's Backhaul rate is the backhaul rate of Isotanks fleet despatched from region Y. To obtain reliable statistics, we perform 100 simulation runs for each scenario and calculate the average and standard deviation of distribution costs, OTIF rates, Fill rates, and Complexity levels.

7.6 The simulation results

We conduct 100 simulation runs for each scenario and calculate the average distribution cost. Table 7.7 shows the calculation results.

Scenario	Region X		Region Y		Region Z		Overall	
	Average Cost	Stdev. ($\sigma_{\bar{x}}$)	Average Cost	Stdev. ($\sigma_{\bar{x}}$)	Average Cost	Stdev. ($\sigma_{\bar{x}}$)	Average Cost	Stdev. ($\sigma_{\bar{x}}$)
X 0% - Y 0%	8,653.2	11.7	4,896.0	9.1	2,659.2	5.3	24,040.6	18.3
X 50% - Y 50%	4,940.3	9.4	4,023.4	7.1	2,662.3	4.9	19,458.3	19.6
X 50% - Y 100%	7,128.0	10.5	3,268.9	6.0	2,663.0	5.9	20,892.1	16.4
X 100% - Y 50%	2,349.9	3.2	4,025.9	6.9	2,659.2	5.0	16,867.2	14.1
X 100% - Y 100%	3,850.1	6.0	3,273.0	6.2	2,662.6	5.4	17,617.9	16.9

Table 7.7: Average cost and Standard Deviation

Based on the average costs in Table 7.7, Figure 7.5 compares the average costs in different scenarios.

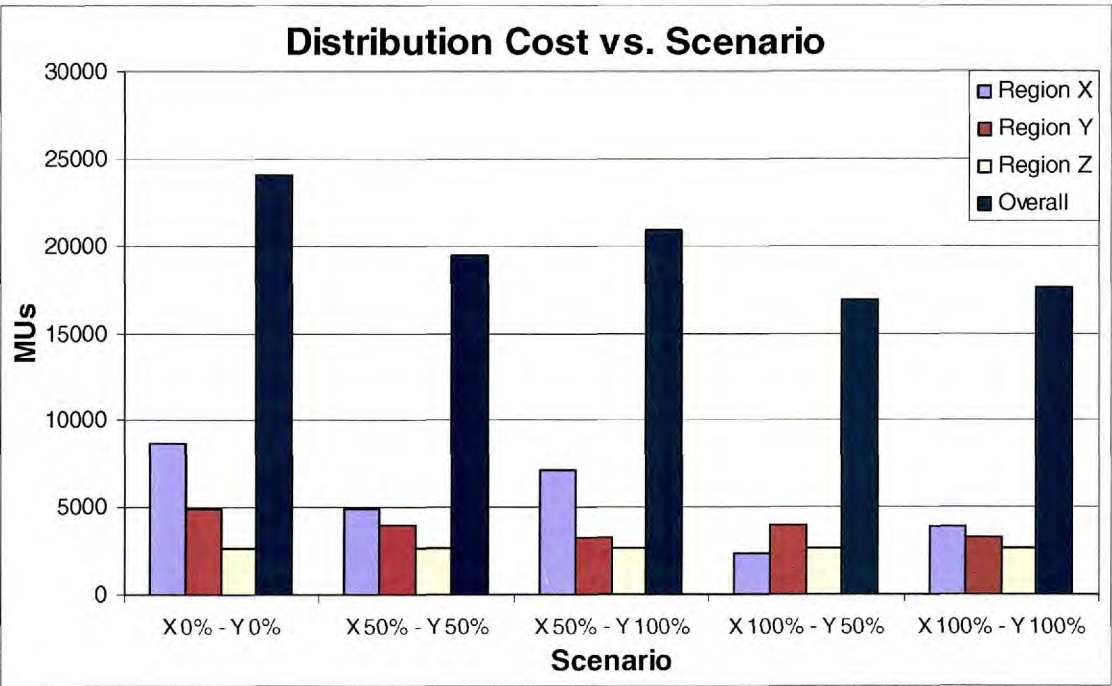


Figure 7.5: Distribution costs in different scenarios

From Figure 7.5, it is evident that the distribution costs decrease as the backhaul strategy is adopted. In general, the distribution costs should further decrease as the backhaul rate is further increased. However, as shown in Figure 7.5, this is not always the case. The distribution costs in the region X Backhaul 50% and region Y Backhaul 100% case are actually higher than those of the region X Backhaul 50% and region Y Backhaul 50% case, even though it has a higher overall backhaul rate. This is because the latter case fails to deliver the products more often than the former case (as revealed by lower OTIF rates and Fill rates in Figures 7.7 and 7.8), hence the distribution costs are lower. The same reason explains the paradox between the region X Backhaul 100% and region Y Backhaul 50% case and the region X Backhaul 100% and region Y Backhaul 100% case.

To avoid the bias due to difference in the amount of product delivered, an alternative approach to compare the distribution costs between different scenarios is to calculate average cost per unit delivered as shown in equation 7.9.

$$\text{Cost per unit delivered} = \frac{\text{Distribution Cost}}{\text{Amount Delivered}} \quad (7.9)$$

The distribution costs in the numerator are calculated using equation 7.1, 7.2, 7.3, and 7.4, respectively. The amount delivered is the amount of product actually delivered from regions X, Y, Z, and all the regions. Table 7.8 summarises the cost per unit delivered for each scenario.

Scenario	Average cost per unit delivered			
	Region X	Region Y	Region Z	Overall
X 0% - Y 0%	0.83	0.87	0.89	1.26
X 50% - Y 50%	1.68	0.66	0.90	1.62
X 50% - Y 100%	0.65	0.54	0.89	1.05
X 100% - Y 50%	1.47	0.66	0.89	1.59
X 100% - Y 100%	1.65	0.54	0.89	1.55

Table 7.8: Average cost per unit delivered in different scenarios

Figure 7.6 compares the cost per unit delivered for the five different scenarios.

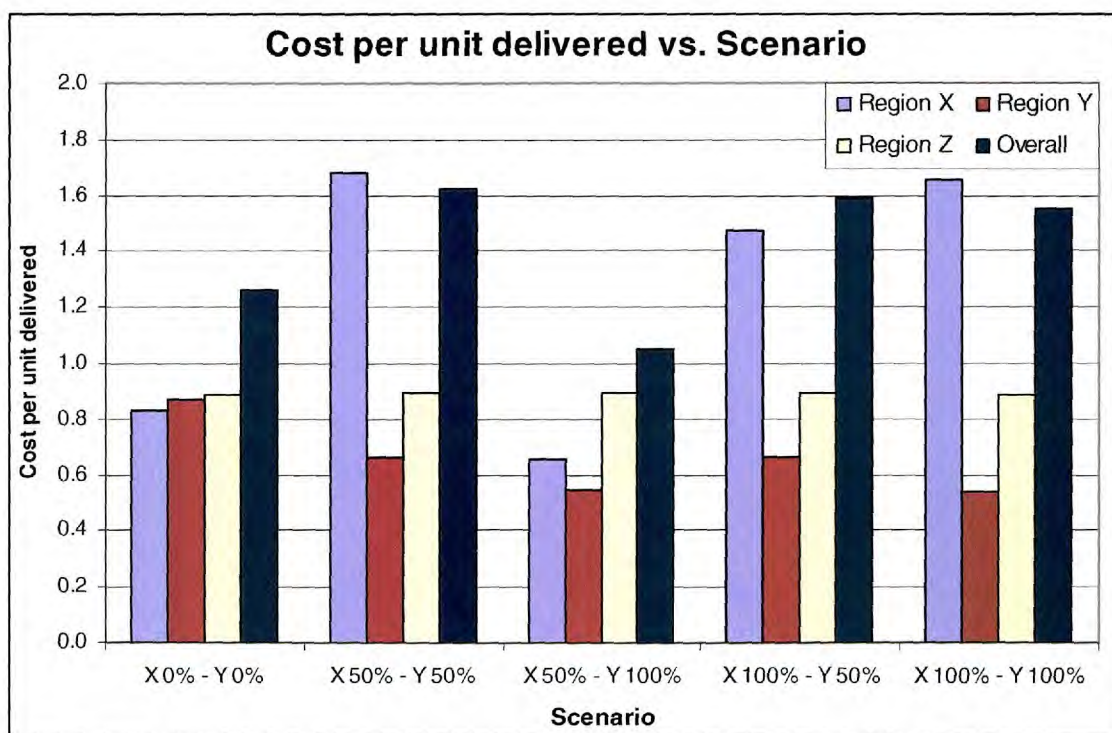


Figure 7.6: Distribution costs per unit delivered in different scenarios

From Table 7.8 and Figure 7.6, it becomes clear that the region X Backhaul 50% and region Y Backhaul 100% case has the lowest distribution cost per unit delivered. Notice that region Y's distribution cost per unit delivered decreases as its backhaul rate increases. However, region X's distribution cost per unit delivered does not always decrease as its backhaul rate increase. This is because it has to keep paying for returning empty Isotanks to region Z (see equation 7.1) and this can be roughly regarded as a fixed cost. Hence, when the amount of product delivered from region X is low (e.g. due to an Isotanks shortage), the distribution cost per unit delivered soars.

The average OTIF rates of the 100 simulation runs are shown in Table 7.9.

Scenario	Region X		Region Y		Region Z		Overall	
	Avg. OTIF rate (%)	Stdev. ($\sigma_{\bar{x}}$)	Avg. OTIF rate (%)	Stdev. ($\sigma_{\bar{x}}$)	Avg. OTIF rate (%)	Stdev. ($\sigma_{\bar{x}}$)	Avg. OTIF rate (%)	Stdev. ($\sigma_{\bar{x}}$)
X 0% - Y 0%	84.83	0.70	82.37	0.86	93.90	0.71	87.03	0.50
X 50% - Y 50%	26.37	0.21	100.00	0.00	93.87	0.65	73.41	0.23
X 50% - Y 100%	98.71	0.39	99.54	0.22	93.90	0.66	97.38	0.28
X 100% - Y 50%	16.58	0.11	100.00	0.00	93.98	0.61	70.19	0.21
X 100% - Y 100%	22.52	0.17	100.00	0.00	94.10	0.75	72.21	0.26

Table 7.9: Average OTIF rate and Standard Deviation

Figure 7.7 compares the average OTIF rates among the different scenarios.

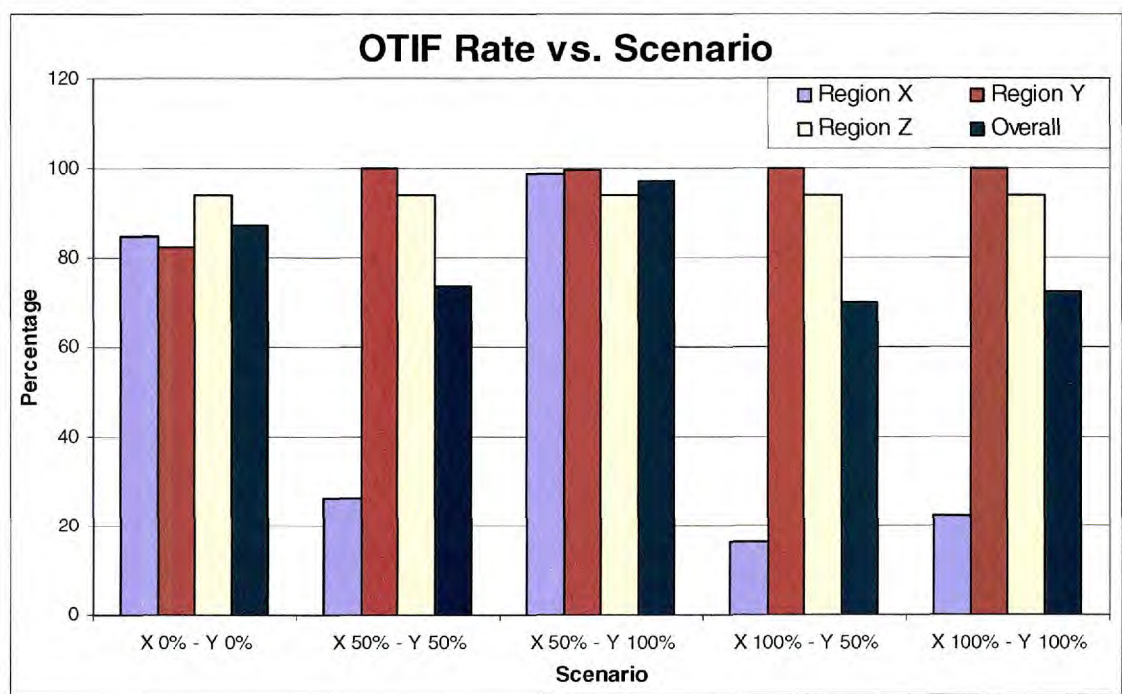


Figure 7.7: OTIF rates in different scenarios

Similarly, the average Fill rates of the 100 simulation runs are shown in Table 7.10.

Scenario	Region X		Region Y		Region Z		Overall	
	Avg. Fill rate (%)	Stdev. ($\sigma_{\bar{x}}$)	Avg. Fill rate (%)	Stdev. ($\sigma_{\bar{x}}$)	Avg. Fill rate (%)	Stdev. ($\sigma_{\bar{x}}$)	Avg. Fill rate (%)	Stdev. ($\sigma_{\bar{x}}$)
X 0% - Y 0%	94.91	0.42	92.67	0.64	98.11	0.33	94.69	0.30
X 50% - Y 50%	26.70	0.17	100.00	0.00	98.23	0.28	59.59	0.12
X 50% - Y 100%	99.57	0.19	99.75	0.12	98.49	0.22	99.46	0.12
X 100% - Y 50%	14.58	0.05	100.00	0.00	98.56	0.24	53.08	0.07
X 100% - Y 100%	21.13	0.11	100.00	0.00	98.38	0.28	56.60	0.09

Table 7.10: Average Fill rate and Standard Deviation

Figure 7.8 compares the average Fill rates among the five different scenarios.

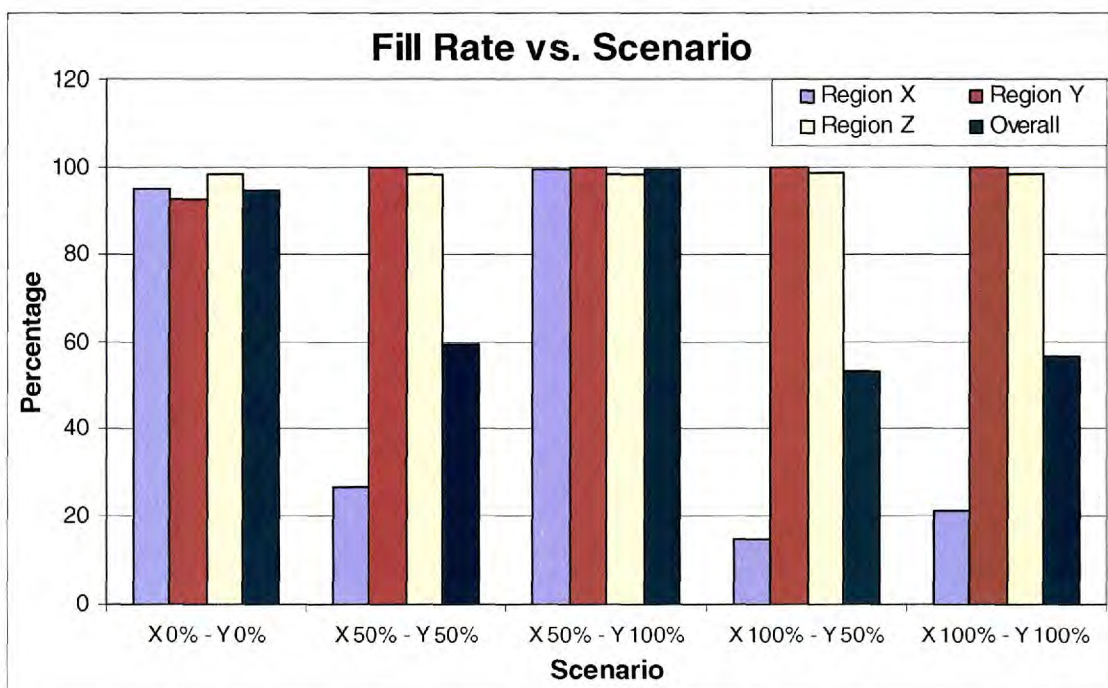


Figure 7.8: Fill rate in different scenarios

In general, the region X Backhaul rate 50% and region Y Backhaul rate 100% case provide the best OTIF rates and Fill rates. Notice that this case also incurs lower distribution costs than the no backhaul case. The OTIF rates and Fill rates of region Z are not affected by the backhaul policy since all Isotanks despatched from region Z will not be backhailed. It is better to do less backhaul for the Isotanks despatched from region X than those from region Y because the product B demand in region Y is greater than product A demand in region X. Therefore, the region X facility will need to keep more Isotanks both from those returned empty from region Y and from those loaded with product A from region Y. On the other hand, the region Y facility requires less Isotanks than region X facility, hence keeping only 50% of Isotanks loaded with product B from region X is sufficient to support the distribution requirements.

Table 7.11 shows average complexity level of each scenario and its standard deviation.

Scenario	Avg. Complexity	Stdev. ($\sigma_{\bar{x}}$)
X 0% - Y 0%	5.301	0.010
X 50% - Y 50%	4.731	0.006
X 50% - Y 100%	4.687	0.009
X 100% - Y 50%	4.396	0.011
X 100% - Y 100%	4.256	0.011

Table 7.11: Average Complexity level and Standard Deviation

Notice that the from 100 simulation runs, the standard deviation (of the average value) becomes very small compared to the average value. Hence, the results are reliable. Figure 7.9 compares the average complexity level for the five different scenarios.

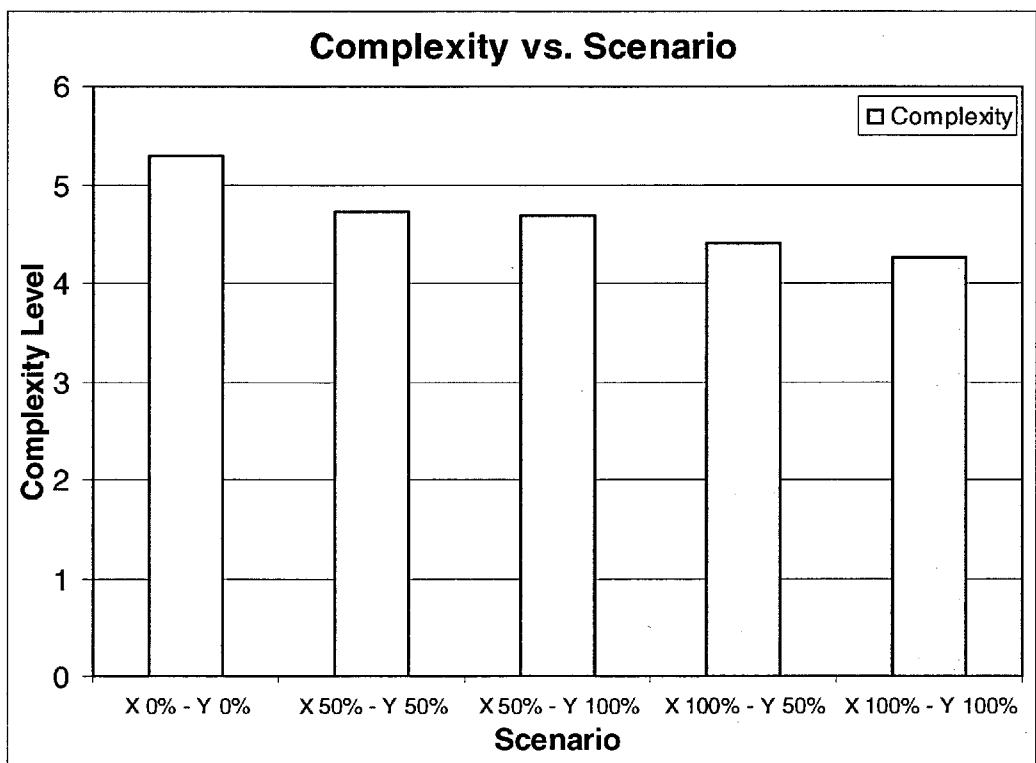


Figure 7.9: Complexity levels in different scenarios

From Figure 7.9, as expected, we can see that the no backhaul case (region X Backhaul rate 0% and region Y Backhaul rate 0%) has the highest level of complexity and the complexity decreases as the backhaul rates increase.

From all the performance measures above, it is evident that the best strategy is to reload all loaded Isotanks from region Y to region X with product B according to product B's demand, otherwise keep the Isotanks at the region X facility, whilst the similar procedure should be done to only half of the loaded Isotanks from region X to region Y and the rest should be returned empty to region X. This will not only improve service rates and reduce distribution cost, but also reduce the complexity of the distribution operation compared to not doing any backhauling at all. Other scenarios may reduce the distribution cost or complexity further, however, the service rates will be unacceptably low.

7.7 Summary and Discussion

The simulation model is proven to be very useful as a decision support tool. It reveals distribution system performance according to the demand profile as different distribution policies are adopted. Distribution costs, service rates, and complexity levels are often conflicting performance measures. Improving service rates usually results in higher distribution costs and potentially a more complex distribution system is required. However, the simulation model shows that this trade-off is not always the case. When the backhaul rates are optimised, distribution costs can be reduced, service rates (OTIF rate and Fill rate) can be improved, whilst the complexity level is lower; this is a very interesting finding. The power of the simulation tool is that it can sometimes reveal potential benefits that may seem counterintuitive. The model user can evaluate as many scenarios as he or she wishes. However, in general, the following principles apply:

1. To reduce distribution costs, a certain level of backhauling should be adopted (to reduce transportation cost) whilst keeping the number of Isotanks as low as possible (to reduce Isotank leasing cost).
2. To maintain high service rates, either do not backhaul at all (however, the service rates will deteriorate as the travel times become longer) or keep a high rate of backhauling for the Isotanks from the origin with lower demand to the destination with higher demand whilst keeping a lower rate of backhauling for the Isotanks from the origin of higher demand to the destination of lower demand.
3. To lower complexity level, try to standardise fleet size, product type, and transportation route or, in other words, try to lower the possibility of having many different states in the system.

However, there are also other strategies that possibly help to further reduce distribution costs and improve service rates. For example, in Thailand, the Tesco supermarket and Tipco foods group both noticed that they have similar vehicle routes but in opposite directions, and often the vehicles return empty to their warehouses. Therefore, they made an agreement to share the vehicles, and hence share the distribution cost which resulted in distribution cost reduction for both companies. However, in this case, the complexity level of the distribution system is likely to increase. Consequently, both companies needed to prepare to cope with a more complex distribution operation, possibly by having effective and efficient IT and transactional systems.

Chapter 8

Conclusions and Further Work

In this thesis, we propose the concept of supply chain complexity and robustness. We argue that supply chain managers should pay attention to these supply chain performance measures since highly complex supply chains will require more resources and effort to synchronise and coordinate supply chain activities and the robustness indicates long-term serviceability of the supply chain. Higher complexity also implies a higher risk of mishandling in managing supply chain. The concept can be applied to any level of supply chain management. In various aspects of supply chain management, we define complexity and robustness measures and the methods to quantify them. In certain supply chain operations the robustness is equivalent to service level.

Through both hypothetical cases and cases from real-world industrial settings, this thesis defines and quantifies complexity and robustness measures that are meaningful for making business decisions. We also derive sensitivity, trade-offs, and relationships among efficiency, complexity, and robustness. Supply chain configurations and vehicle routes can also be regarded as networks of nodes and edges, hence, we relate them to graph theory and analyse the networks' topological properties. In addition, this thesis also presents improvements of demand forecasting accuracy through quantitative forecasting methods. The work in this thesis is highly relevant and adds value to supply chain researchers and practitioners. Through both hypothetical and real cases, we obtain several original and useful results and relationships which can be concluded as in the next section.

8.1 Conclusions

To demonstrate our concept, in chapter 3, we start from one of the earliest strategic decisions in supply chain management, the logistical network design, by developing a framework and a mathematical model to optimise efficiency (cost) whilst satisfying different levels of robustness requirements. Logistical network complexity is defined as the ratio of actual number of network edges to the number of minimum spanning tree edges of the smallest network of the same problem setting. From the logistical networks obtained, we measure logistical network complexity and identify its relationships with robustness and efficiency. We also calculate average path length, clustering coefficient, and construct degree distribution of the resulting network topologies based on a graph theoretic formalism. The results are then compared with those of random networks of the same size and average degree. The results show that we can obtain logistical network configurations with desirable robustness levels whilst minimising cost. We also present the relationships of robustness versus normalised efficiency and complexity. The results also show that relationships between logistical network topological properties and robustness exist, as in other real world natural and man-made complex networks. The results from our logistical network design framework were published in Meepetchdee and Shah (2007a).

In chapter 4, within the network configurations resulted from the logistical network design, we simulate vehicle routing and measure complexity and topological properties of different vehicle routing networks with different demand profiles, vehicle sizes, and time horizons. In this chapter, we adopt the entropic complexity measure approach as it reflects the expected amount of information required to describe the state of the vehicle routing operations which implies the amount of coordination and planning effort. We also construct the degree distributions of vehicle routing networks and calculate the statistical measures of the degree of vertex. We also explore the relationship between complexity level and the number of trucks required for vehicle routing. The results show that centralised distribution and low vehicle capacities tend to result in higher complexity levels. We also found both power-law and non power-law degree distributions for vehicle routing networks. In general, vehicle routing networks with high vehicle capacities and shorter planning horizons exhibit approximate power-law distributions and scale-free characteristics. Our approach provides insights into complexity in vehicle routing operations which will help supply chain managers to identify appropriate logistics configurations and/or vehicle capacities, etc. The results from our vehicle routing complexity analysis were submitted for publication in Meepetchdee and Shah (2007b).

Business entities within the supply chain nowadays are becoming more and more distributed in terms of their decision making processes. Nonetheless, every member in the supply chain is still interdependent and needs to interact with other members within the supply chain to achieve their goals. To gain valid insights from supply chain systems analysis, it is essential that the tools used can effectively capture and characterise the discrete, dynamic, and distributed nature of such supply chain systems. In chapter 5, we propose a multi-agent based simulation as a tool to analyse the performance of such a distributed supply chain system. As it is widely documented that an IT revolution has taken place in supply chain management, we use this simulation technique to assess the benefits of potentially more complex transactional systems. We create a simulation model on the NetLogo modelling platform and perform simulation on our hypothetical case study to evaluate supply chain performance in supply chain environments under different conditions of operational complexity. The results of the simulation show that transactional degrees of freedom such as fixed order increments affect complexity level as well as certain supply chain performance measures. We can also observe several useful relationships between supply chain performance measures. For example, fixed order increments strongly relate to inventory and transactional complexity in the supply chain. Although a higher level of complexity is normally undesirable, compromising on a higher level of complexity can be beneficial. Our results reveal the benefits, particularly increasing inventory turnover and lowering average number of orders, gained from accepting higher level of transactional complexity. The results such as the trade-off between inventory management complexity (MAPE) and transactional complexity and between inventory turnover ratio and average number of orders are very revealing and useful for supply chain managers. In general, these behaviours are generally known as 'emergent behaviours' which are the outcomes of each agent's behaviours and their interaction with other agents in the supply chain. The results from our multi-agent based supply chain simulation were submitted for publication in Meepetchdee and Shah (2007c).

Forecasting also plays a central role in the supply chain operations of a firm since all business require planning based on forecasts. Chapter 6 discusses the application of a quantitative forecasting technique as a tool to forecast future product demand at INEOS Fluor Ltd. where only qualitative forecasting methods have been used. In a collaborative project with INEOS Fluor, we do not only provide efficient and more accurate demand forecasting tools but also provide a framework to evaluate forecasting errors and analyse residual errors. By comparison to the traditional practice, when appropriately selected, quantitative forecasting methods provide significant improvements in forecasting accuracy. Quantitative forecasting results can also be used as a benchmark or a reality check for INEOS Fluor's traditional forecasts.

Backhauling is a product distribution strategy that can open up a cost reduction opportunity for INEOS Fluor. However, due to the demand mismatch between different markets, full backhauling can cause Isotank shortages at one location whilst having excessive Isotanks at the other location. In chapter 7, we propose a multi-agent based simulation tool which can reveal potential trade-off among efficiency (cost), robustness (service levels), and system complexity. The simulation model takes several inputs including demand forecasts such as those discussed in chapter 6. One of the advantages of the simulation tool is that it can sometimes reveal potential benefits that may seem counterintuitive. Based on our scenario analysis via the simulation, when the backhaul rate is appropriately chosen, the company can not only reduce distribution cost but also improve service rates and lower the complexity level.

This thesis brings together a number of available operations research tools, management science tools, and econometric tools to design, analyse, and evaluate different aspects of supply chain system in different circumstances under our newly developed framework of efficiency, complexity, and robustness considerations. This framework integrates the considerations for profitability (efficiency), serviceability (robustness), and hidden risk (complexity) of the supply chain system. It provides a more comprehensive view of supply chain system and shall complement all other supply chain analysis frameworks both in the literature and in practice.

8.2 Recommendations for Future Work

This thesis conceptually demonstrates a framework for designing and analysing a supply chain system with considerations of efficiency, complexity, and robustness. We explore this framework in logistical network design, vehicle routing, delivery transaction, and distribution planning with both hypothetical cases and cases from real industrial settings. We also investigate and develop a quantitative forecasting tool to improve demand forecasting as commissioned by the management of INEOS Fluor Ltd. However, supply chain management is a rapidly evolving discipline which requires continuous improvement. This thesis accomplishes some work that contributes to the development of this field. Nonetheless, there is still room for further research and the possible directions for future work are:

Investigating a more complex supply chain system

It is interesting to extend the framework beyond the scale of the supply chain systems addressed in this work. Most commercial supply chains involve a large number of products, span several supply chain echelons, and are horizontally distributed. The complexity normally increases as

more echelons are taken into considerations and when supply chains become more distributed. However, such supply chains also offer greater opportunity to identify non-value added complexity and to reduce complexity level further. Uncertainty and interdependency in a supply chain and its environment are also commonplace and can be taken into account when evaluating or designing the supply chain. Complexity can also cause “nervousness” in the system, for example in production scheduling, a highly complex production flow may result in frequent changes in the production schedule. This change can propagate into other echelons in the supply chain such as in material suppliers or customers. One of the most well-known system behaviours that propagate down the supply chain is the Bullwhip effect. It is interesting to investigate the relationship between the Bullwhip effect and complexity across the multiple layers in a supply chain, potentially through a simulation technique. Such an investigation will provide insights to support both corporate level decisions and operational level decisions.

Linking complexity and robustness to financial benefit or burden

In this thesis, we make no attempt to directly quantify the financial benefit or burden of complexity and robustness. In some work such as Statistical Inventory Models, cost of good will and stock out cost can be explicitly defined to develop the optimal inventory management policy. By the same token, once the cost or benefit of complexity and robustness are defined, we can incorporate them into the objective function of the optimisation model. From an enterprise-wide perspective, it might be desirable to transfer complexity from one member in the supply chain to another since the cost of complexity in one operation may be higher or lower than the others. By formulating a corresponding optimisation model, a supply chain manager will be able to decide where the complexity should be transferred to or taken from. For example, Hewlett-Packard (HP) adopted a postponement strategy to postpone packaging and customisation processes of their DeskJet printers for the European markets (Feitzinger and Lee, 1997). Under the postponement strategy, HP ships uncustomised printers from the US to their distribution centres in Europe where the printers will be customised and packaged. The advantage of such a strategy is that performing the customisation process in Europe provides better responsiveness to changes in customers’ demand and preference, whilst the transportation time from the US to Europe is long, and keeping uncustomised printers results in a lower inventory level due to the risk pooling effect. However, effectively this strategy transfers complexity from the US manufacturing sites to the European distribution centres. Although the postponement strategy is qualitative in nature, it would be valuable to approach the strategy from a quantitative framework by explicitly quantifying cost of complexity and benefit of postponement to determine to what extent such operations should be postponed, potentially via an optimisation model.

Multi-objective supply chain optimisation

In this thesis, robustness is treated as a constraint whilst complexity is not yet incorporated directly into the optimisation model. As a robust supply chain is more desirable, it is logical that we should also attempt to maximise supply chain robustness. On the contrary, a supply chain manager might want to minimise supply chain complexity. This gives rise to a multi-objective optimisation problem. Several techniques for modelling and solving multi-objective optimisation problems are presented in the literature. It is beneficial to apply the available techniques and/or to identify a new modelling approach to optimise efficiency, complexity, and robustness simultaneously.

Coupling optimisation and simulation tools

Whilst optimisation techniques can yield good or the best possible solution under particular circumstance and objectives, they normally deal with static models and do not take into account changes over time. However, supply chain infrastructure and environment are subject to constant change. The optimisation technique alone may not be adequate as a decision support tool for certain problem settings. On the other hand, simulation techniques can take into account the detailed dynamic behaviour of the system and is frequently an effective tool to help analysing problem with some random, or stochastic, elements. In many real-world situations, each member in the supply chain locally optimises their operations upon available information and resources whilst interacting with other members in the supply chain. To evaluate such supply chain behaviour, there is a need to combine the optimisation tool with a simulation tool. For example, Gjerdrum *et al.* (2001b) presented a combined optimisation and agent-based approach to supply chain modelling and performance assessment. In chapter 5, we may extend the work further by adding decision agents who optimise and plan production schedules and product distribution whilst performing a simulation. It is valuable to investigate and compare the efficiency, complexity, and robustness of such a system when different policies are adopted, for example between decentralised optimisation and centralised optimisation.

Explanatory forecasting models

Lastly, there is still room for further improvement regarding demand forecasting. The exponential smoothing techniques and the time series decomposition methods generally capture certain statistically identifiable patterns in historical data without trying to establish any explanation or to identify the factors which influence such patterns. These two techniques can

be quickly developed in the early stages of adopting quantitative forecasting practices with a reasonable level of accuracy, especially when historical data exhibit trend or patterns. However, although significantly more time and effort will be required, it is worthwhile to develop an explanatory forecasting model. Although an explanatory model does not necessarily improve forecasting accuracy, it can provide insights into which factors can effect demand and offer an opportunity for management to manipulate demand to achieve their desirable level. In addition, our work in this thesis is to develop forecasting models for aggregate demand at regional level. It would also be beneficial to disaggregate and further develop forecasting models for demand at individual customer level or at national level.

Supply chain environments and practices change constantly, and as continuous improvement is a necessity in supply chain management, supply chain research activities must be able to respond and address these issues in a timely manner.

References

- Ahn, H. J., Park, S. J. (2003), Modeling of a Multi-agent System for Coordination of Supply Chains with Complexity and Uncertainty, In *Intelligent Agents and Multi-Agent Systems*, edited by J. Lee and M. Barley, pp. 13-24, Springer-Verlag, Germany
- Albert, R., Barabási, A.L. (2000) Topology of Evolving Networks: Local Events and Universality, *Physical Review Letters* 85 (24), 5234-5237
- Albert, R., Barabási, A.L. (2002), Statistical mechanics of complex networks, *Reviews of Modern Physics* 74 (1), 47-97
- Albert, R., Jeong, H., Barabási, A.L. (1999), Diameter of the World-Wide Web, *Nature* 401 (6749), 130-131
- Allwood, J. M., Lee, J. H. (2005), The design of an agent for modelling supply chain network dynamics, *International Journal of Production Research* 43 (22), 4875-4898
- Anon. (2006), £6 billion shopping spree in two days, *METRO*, 20 December 2006
- Archetti, C., Mansini, R., Speranza, M. G. (2005), Complexity and Reducibility of the Skip Delivery Problem, *Transportation Science* 39 (2), 182-187
- Bailar, B. A., Lanphier, M. C. (1978), *Development of Survey Methods to Assess Survey Practices*, American Statistical Association, Washington, USA
- Ballou, R. H. (2001), Unresolved Issues in Supply Chain Network Design, *Information Systems Frontiers* 3 (4), 417-426

- Barabási, A.L., Jeong, H., Néda, Z., Ravasz, E., Schubert, A., Vicsek, T. (2002), Evolution of the social network of scientific collaborations, *Physica A: Statistical Mechanics and its Applications* 311 (3-4), 590-614
- Bender, P. S. (2000), Debunking 5 Supply Chain Myths, *Supply Chain Management Review* 4 (1), 52-58
- Berger, J., Barkaoui, M. (2004), A parallel hybrid genetic algorithm for the vehicle routing problem with time windows, *Computers & Operations Research* 31 (12), 2037-2053
- Bernstein, P. L. (1996), *Against The Gods: The Remarkable Story of Risk*, John Wiley & Sons, Chichester, UK
- Bose, S., Pekny, J. F. (2000), A model predictive framework for planning and scheduling problems: a case study of consumer goods supply chain, *Computers & Chemical Engineering* 24 (2-7), 329-335
- Box, G. E. P., Jenkins, G. M. (1970), *Time Series Analysis, Forecasting, and Control*, Holden Day, San Francisco, USA
- Calinescu A., Bermejo, J., Efstathiou, J., Schirn, J. (1997), Applying and Assessing Techniques For Measuring Complexity In Manufacturing, In Proceedings of the First European Conference on Intelligent Management Systems in Operations, 25-26 March 1997, University of Salford, Greater Manchester, UK, pp. 21-28
- Calinescu, A., Efstathiou, J., Schirn, J., Bermejo, J. (1998), Applying and Assessing Two Methods for Measuring Complexity in Manufacturing, *The Journal of the Operational Research Society* 49 (7), 723-733
- Calinescu, A., Efstathiou, J., Sivadasan, S., Schirn, J., Huaccho L. H. (2000), Complexity in manufacturing: An information theoretic approach, In Proceedings of "Complexity and Complex Systems in Industry", 18-20 September 2000, Eds. McCarthy IP, Rakatobe-Joel T, University of Warwick, pp. 30-44
- Caridi, M., Cavalieri, S. (2004), Multi-agent systems in production planning and control: an overview, *Production Planning & Control* 15 (2), 106-118

- Chambers, J. C., Mullick, S. K., Smith, D. D. (1971), How to choose the right forecasting technique, *Harvard Business Review* 49 (4), 45-74
- Chepuri, K., Homem-De-Mello, T. (2005), Solving the Vehicle Routing Problem with Stochastic Demands using the Cross-Entropy Method, *Annals of Operations Research* 134 (1), 153-181
- Choi, H. R., Kim, H. S., Park, B. J., Park, Y. S. (2004), Multi-agent Based Integration Scheduling System under Supply Chain Management Environment, In *Innovations in Applied Artificial Intelligence*, edited by B. Orchard, C. Yang, and M. Ali, pp. 249-263, Springer-Verlag, Germany
- Choi, T. Y., Krause, D. R. (2006), The supply base and its complexity: Implications for transaction costs, risks, responsiveness, and innovation, *Journal of Operations Management* 24 (5), 637-652
- Chopra, S. (2003), Designing the distribution network in a supply chain, *Transportation Research Part E: Logistics and Transportation Review* 39 (2), 123-140
- Chopra, S., Meindl, P. (2004), *Supply Chain Management: Strategy, Planning, and Operations*, 2nd ed., Pearson Prentice Hall, New Jersey, USA
- Chopra, S., Sodhi, M. (2004), Managing risk to avoid supply-chain breakdown, *Sloan Management Review* 46 (1), 53-62
- Christopher, M. (2004), Creating resilient supply chains, *Logistics Europe*, Feb., 14-21
- Christopher, T. S. (2006), Robust strategies for mitigating supply chain disruptions, *International Journal of Logistics: Research and Applications* 9 (1), 33-45
- Cohen, M. A., Lee, H. L. (1989), Resource Deployment Analysis of Global Manufacturing and Distribution Networks, *Journal of Manufacturing and Operations Management* 2, 81-104
- Cohen, M. A., Moon, S. (1990), Impact of Production Scale Economies, Manufacturing Complexity, and Transportation Costs on Supply Chain Facility Networks, *Journal of Manufacturing and Operations Management* 3, 269-292

- Czech, Z., Czarnas, P. (2002), Parallel simulated annealing for the vehicle routing problem with time windows, In Proceedings of the 10th Euromicro Workshop on Parallel, Distributed and Network-based Processing, Spain, pp.376-383
- Dalrymple, D. J. (1987), Sales Forecasting Practices: Results from a United States Survey, *International Journal of Forecasting* 3 (3-4), 379-391
- Davidsson, P., Wernstedt, F. (2004), A Framework for Evaluation of Multi-Agent System Approaches to Logistics Network Management, In *An Application Science for Multi-Agent Systems*, edited by T.A. Wagner, pp. 27-40, Kluwer Academic Publishers, USA
- Deo, N. (1974), *Graph Theory with Applications to Engineering and Computer Science*. Prentice-Hall series in automatic computation, Prentice-Hall, New Jersey, USA
- de Paepe, W. E., Lenstra, J. K., Sgall, J., Sitters, R. A., Stougie, L. (2004), Computer-Aided Complexity Classification of Dial-a-Ride Problems, *INFORMS Journal on Computing* 16 (2), 120-132
- Deshmukh, A. V. (1993), *Complexity and chaos in manufacturing systems*, Ph.D. thesis, Purdue University, USA
- Efstathiou, J., Kariuki, S., Huatuco L. H., Sivadasan, S., Calinescu, A. (2001), The relationship between information-theoretic and chaos-theoretic measures of the complexity of manufacturing systems, In Proceedings of the 17th National Conference on Manufacturing Research (NCMR 2001): Advances in Manufacturing Technology XV. 4th– 6th September 2001, University of Cardiff, UK, pp. 421-426
- Eibl, P. G., Mackenzie, R., Kidner, D. B. (1994), Vehicle Routeing and Scheduling in the Brewing Industry: A Case Study, *International Journal of Physical Distribution & Logistics Management* 24 (6), 27-37
- Feitzinger, E., Lee, H. L. (1997), Mass Customization at Hewlett-Packard: The Power of Postponement, *Harvard Business Review* 75 (1), 116-121
- Ferber, J. (1999), *Multi-Agent Systems: An Introduction to Distributed Artificial Intelligence*, Addison-Wesley, Singapore

- Fischer, T., Gehring, H. (2005), Planning vehicle transshipment in a seaport automobile terminal using a multi-agent system, *European Journal of Operational Research* 166 (3), 726-740
- Fisher, M. L., Hammond, J. H., Obermeyer, W. R., Raman, A. (1994), Making Supply Meet Demand in an Uncertain World, *Harvard Business Review* 72 (3), 83-93
- Franses, P. H. (2004), Do We Think We Make Better Forecasts Than in the Past? A Survey of Academics, *Interfaces* 34 (6), 466-468
- Frizelle, G., Woodcock, E. (1995), Measuring complexity as an aid to developing operational strategy, *International Journal of Operations & Production Management* 15 (5), 26-39
- Fu, M. C. (2002), Optimization for Simulation: Theory vs. Practice, *INFORMS Journal on Computing* 14 (3), 192-215
- Funk, J. L. (1995), Just-in-time manufacturing and logistical complexity: a contingency model, *International Journal of Operations & Production Management* 15 (5), 60-71
- García-Flores, R., Wang, X. Z. (2002), A multi-agent system for chemical supply chain simulation and management support, *OR Spectrum* 24 (3), 343-370
- Gino, F. (2002), Complexity measures in decomposable structures, Working Paper prepared for the 2nd EURAM (European Academy of Management) Conference on "Innovative Research in Management", 9-11 May 2002, Stockholm, Sweden, Available online at: http://www.ecsocman.edu.ru/images/pubs/2002/12/11/0000016270/complexity_measures.pdf (accessed 9 April 2007)
- Gjerdrum, J., Shah, N., Papageorgiou, L. G. (2001a), Transfer Prices for Multienterprise Supply Chain Optimization, *Industrial and Engineering Chemistry Research* 40 (7), 1650-1660
- Gjerdrum, J., Shah, N., Papageorgiou, L. G. (2001b), A combined optimization and agent-based approach to supply chain modelling and performance assessment, *Production Planning & Control* 12 (1), 81-88
- Guimaraes, T., Martensson, N., Stahre, J., Igbaria, M. (1999), Empirically testing the impact of manufacturing system complexity on performance, *International Journal of Operations and Production Management* 19 (12), 1254-1269

- Harrison, T. P. (2001), Global Supply Chain Design, *Information Systems Frontiers* 3 (4), 413-416
- Hassin, R., Rubinstein, S. (2005), On the complexity of the k-customer vehicle routing problem, *Operations Research Letters* 33 (1), 71-76
- Hax, A. C., Candea, D. (1984), *Production and Inventory Management*, Prentice Hall, New Jersey, USA
- Helmer, O., Rescher, N. (1959), On the Epistemology of the Inexact Sciences, *Management Science*, 6 (1), 25-52
- Hillier, F. S., Lieberman, G. J. (2005), *Introduction to Operations Research*, 8th ed., McGraw-Hill, Singapore
- Hopp, W. J., Spearman, M. L. (2000), *Factory Physics*, 2nd ed., McGraw-Hill, Singapore
- Howarth, S. (1997), *A Century in Oil: The "Shell" Transport and Trading Company 1897-1997*, Weidenfeld & Nicolson, London, UK
- Hurwood, D. L., Grossman, E. S., Bailey, E. L. (1978), *Sales Forecasting*, The Conference Board, New York, USA
- Hwang, H. S. (2002), An improved model for vehicle routing problem with time constraint based on genetic algorithm, *Computers & Industrial Engineering* 42 (2-4), 361-369
- Jalisi, Q. W. Z. (2001), *Logistics and supply chain optimal planning and scheduling*, Ph.D. thesis, Imperial College London, UK
- Jennings, N. R., Woodridge, M. J. (1998), *Agent Technology: Foundations, Applications, and Markets*, Springer-Verlag, Germany
- Jeong, H., Tombor, B., Albert, R., Oltvai, Z. N., Barabási, A.L. (2000), The large-scale organization of metabolic networks, *Nature* 407 (6804), 651-654

- Julka, N., Srinivasan, R., Karimi, I. (2002a), Agent-based supply chain management-1: framework, *Computers and Chemical Engineering* 26 (12), 1755-1769
- Julka, N., Srinivasan, R., Karimi, I. (2002b), Agent-based supply chain management-2: a refinery application, *Computers and Chemical Engineering* 26 (12), 1771-1781
- Jung, J. Y., Blau, G., Pekny, J. F., Reklaitis, G. V., Eversdyk, D. (2004), A simulation based optimization approach to supply chain management under demand uncertainty, *Computers & Chemical Engineering* 28 (10), 2087-2106
- Kaihara, T. (2003), Multi-agent based supply chain modelling with dynamic environment, *International Journal of Production Economics* 85 (2), 263-269
- Kamrad, B., Siddique, A. (2004), Supply Contracts, Profit Sharing, Switching, and Reaction Options, *Management Science* 50 (1), 64-82
- Kauffman, S. (1993), *The Origins of Order*, Oxford University Press, Oxford, UK
- Keilmann, K. P. (1995), Multi-Agent-Systems- A Natural Trend in CIM, In *Information Management in Computer Integrated Manufacturing: A Comprehensive Guide to State-of-the-Art CIM Solutions*, edited by H.H. Adelsberger, J. Lažanský, and V. Mařík, pp. 548-562, Springer-Verlag, Germany
- Lasdon, L. S., Waren, A. D., Jain, A., Ratner, M. (1978), Design and Testing of a Generalized Reduced Gradient Code for Nonlinear Programming, *ACM Transactions on Mathematical Software* 4 (1), 34-50
- Law, A. M., Kelton, W. D. (2000), *Simulation Modeling and Analysis*, 3rd ed., McGraw-Hill, USA
- Leahy, S. E., Murphy, P. R., Poist, R. F. (1995), Determinants of Successful Logistical Relationships: A Third-Party Provider Perspective, *Transportation Journal* 35 (2), 5-13
- Lee, H. L. (2004), The Triple-A Supply Chain, *Harvard Business Review* 82 (10), 102-112
- Lee, H. L., Padmanaphan, V., Wang, S. (1997), Information Distortion in a Supply Chain: The Bullwhip Effect, *Management Science* 43 (4), 546-558

- Lee, H. L., Sasser, M. M. (1995), Product universality and design for supply chain management, *Production Planning & Control* 6 (3), 270-277
- Lee, K., Kim, W., Kim, M. (2004), Supply Chain Management using Multi-agent System, In *Cooperative Information Agents VIII*, edited by M. Klusch, S. Ossowski, V. Kashyap, and R. Unland, pp. 215-225, Springer-Verlag, Germany
- Lee, T. R., Ueng, J. H. (1999), A study of vehicle routing problems with load-balancing, *International Journal of Physical Distribution & Logistics Management* 29 (10), 646-658
- Levin, R. I., Rubin, D. S. (1997), *Statistics for Management*, 7th ed., Prentice-Hall, New Jersey, USA
- Levis, A. A., Papageorgiou, L. G. (2005), Customer Demand Forecasting Via Support Vector Regression Analysis, *Chemical Engineering Research and Design*, 83 (A8), 1009-1018
- Li H., Lim, A. (2003), Local search with annealing-like restarts to solve the VRPTW, *European Journal of Operational Research* 150 (1), 115-127
- Lin, F. R., Pai, Y. H. (2000), Using Multi-Agent Simulation and Learning to Design New Business Process, *IEEE Transactions on Systems, Man, and Cybernetics-Part A* 30 (3), 380-384
- Lin, F. R., Shaw, M. J. (1998), Reengineering the Order Fulfillment Process in Supply Chain Networks, *International Journal of Flexible Manufacturing Systems* 10 (3), 197-229
- Lin, F. R., Tan, G. W., Shaw, M. J. (1998), Modeling Supply-Chain Networks by a Multi-Agent System, In *Proceedings of the IEEE 31st Hawaii International Conference on System Sciences*, pp. 105-114
- Makridakis, S., Wheelwright, S. C. (1989), *Forecasting Methods for Management*, 5th ed., John Wiley & Sons, New York, USA
- Meepetchdee, Y., Shah, N. (2007a), Logistical Network Design with Robustness and Complexity Considerations, *International Journal of Physical Distribution & Logistics Management* 37 (3), 201-222

- Meepetchdee, Y., Shah, N. (2007b), Complexity and topological properties of vehicle routing networks, *Transportation Research Part E: Logistics and Transportation Review* [In Review]
- Meepetchdee, Y., Shah, N. (2007c), Supply chain operational performance and complexity assessment through multi-agent based simulation, *International Journal of Production Research* [In Review]
- Mentzer, J. T., Cox, J. E. (1984), Familiarity Application and Performance of Sales Forecasting Techniques, *Journal of Forecasting* 3 (1), 27-36
- Mohamed, Z. M. (1999), An integrated production-distribution model for a multi-national company operating under varying exchange rates, *International Journal of Production Economics* 58 (1), 81-92
- Montané, F. A. T., Galvão, R. D. (2006), A tabu search algorithm for the vehicle routing problem with simultaneous pick-up and delivery service, *Computers & Operations Research* 33 (3), 595-619
- Montoya, J. M., Solé, R. V. (2002), Small World Patterns in Food Webs, *Journal of Theoretical Biology* 214 (3), 405-412
- Muckstadt, J. A., Murray, D. H., Rappold, J. A., Collins, D. E. (2001), Guidelines for Collaborative Supply Chain System Design and Operation, *Information Systems Frontiers* 3 (4), 427-453
- Nahmias, S. (1994), Demand Estimation in Lost Sales Inventory Systems, *Naval Research Logistics* 41 (6), 739-757
- Nahmias, S. (2005), *Production and Operations Analysis*, 5th ed., McGraw-Hill, Singapore
- Narus, J. A., Anderson, J. C. (1986), Turn your industrial distributors into partners, *Harvard Business Review* 64 (2), 66-71
- Pankratz, G. (2005a), A Grouping Genetic Algorithm for the Pickup and Delivery Problem with Time Windows, *OR Spectrum* 27 (1), 21-41
- Pankratz, G. (2005b), Dynamic vehicle routing by means of a genetic algorithm, *International Journal of Physical Distribution & Logistics Management* 35 (5), 362-383

- Parunak, H. V. D. (1997), "Go to the ant": Engineering principles from natural multi-agent systems, *Annals of Operations Research* 75, 69-101
- Pěchouček, M., Vokřínek, J., Hodík, J., Bečvář, P., Pospíšil, J. (2004), ExPlanTech: Multi-agent Framework for Production Planning, Simulation and Supply Chain Management, In *Multiagent System Technologies*, edited by G. Lindemann, J. Denzinger, I.J. Timm, and R. Unland, pp. 287-300, Springer-Verlag, Germany
- Phukan, A., Kalava, M., Prabhu, V. (2005), Complexity metrics for manufacturing control architectures based on software and information flow, *Computers & Industrial Engineering* 49 (1), 1-20
- Porter, M. E. (1980), *Competitive Strategy*, The Free Press, New York, USA
- Ropke, S., Pisinger, D. (2006), A unified heuristic for a large class of Vehicle Routing Problems with Backhauls, *European Journal of Operational Research* 171 (3), 750-775
- Rowley, I. (2005), Japan Isn't Buying Wal-Mart Idea, *BusinessWeek*, 28 February 2005
- Sanders, N. R., Manrodt, K. B. (2003), Forecasting Software in Practice: Use, Satisfaction, and Performance, *Interfaces* 33 (5), 90-93
- Sateesh, C., Ray, P. K. (1992), Vehicle Routeing in Large Organizations: A Case Study, *International Journal of Operations & Production Management* 12 (3), 71-78
- Satin, A., Shastry, W. (1983), *Survey Sampling: A Non-mathematical Guide*, Statistics Canada, Ottawa, Canada
- Schonberger, R. J. (1996), Strategic Collaboration: Breaching the Castle Walls, *Business Horizons* 39 (2), 20-26
- Shannon, C. E. (1948), The Mathematical Theory of Communication, *The Bell System Technical Journal* 27 (3), 379-423
- Shapiro, J. F. (2001), *Modeling the Supply Chain*, Duxbury, USA

- Shapiro, J. F. (2004), Challenges of strategic supply chain planning and modelling, *Computers & Chemical Engineering* 28 (6-7), 855-861
- Sheffi, Y. (2001), Supply Chain Management under the Threat of International Terrorism, *International Journal of Logistics Management* 12 (2), 1-11
- Simchi-Levi, D., Kaminsky, P., Simchi-Levi, E. (2003), *Designing & Managing the Supply Chain: Concepts, Strategies & Case Studies*, 2nd ed., McGraw-Hill, Singapore
- Sivadasan, S., Efstathiou, J., Frizelle, G., Shirazi, R., Calinescu, A. (2002), An information-theoretic methodology for measuring the operational complexity of supplier-customer systems, *International Journal of Operations & Production Management* 22 (1), 80-102
- Swaminathan, J. M. (2001), Enabling Customization Using Standardized Operations, *California Management Review* 43 (3), 125-135
- Swaminathan, J. M., Smith, S. F., Sadeh, N. M. (1998), Modeling Supply Chain Dynamics: A Multiagent Approach, *Decision Sciences* 29 (3), 607-632
- Tan, K. C., Lee, L. H., Zhu, Q. L., Ou, K. (2001), Heuristic methods for vehicle routing problem with time windows, *Artificial Intelligence in Engineering* 15 (3), 281-295
- Tryfos, P. (1996), *Sampling Methods for Applied Research: Text and Cases*, Wiley, New York, USA
- Tsiakis, P., Shah, N., Pantelides, C. C. (2001), Design of Multi-echelon Supply Chain Networks under Demand Uncertainty, *Industrial and Engineering Chemistry Research* 40 (16), 3585-3604
- van der Zee, D. J., van der Vorst, J. G. A. J. (2005), A Modeling Framework for Supply Chain Simulation: Opportunities for Improved Decision Making, *Decision Sciences* 36 (1), 65-95
- van Hoek, R. I., Vos, B., Commandeur, H. R. (1999), Restructuring European Supply Chains by Implementing Postponement Strategies, *Long Range Planning* 32 (5), 505-518

- Venkatasubramanian, V., Katare, S., Patkar, P. R., Mu, F. (2004), Spontaneous emergence of complex optimal networks through evolutionary adaptation, *Computers & Chemical Engineering* 28 (9), 1789-1798
- Wassan, N. A. (2006), A reactive tabu search for the vehicle routing problem, *Journal of the Operational Research Society* 57 (1), 111-116
- Watts, D. J., Strogatz, S. H. (1998), Collective dynamics of 'small-world' networks, *Nature* 393 (6684), 440-442
- Wilensky, U. (1999), NetLogo, Center for Connected Learning and Computer-based Modeling, Northwestern University, USA, Available online at:
<http://ccl.northwestern.edu/netlogo/> (accessed 9 April 2007)
- Wooldridge, M. (2002), *An Introduction to MultiAgent Systems*, John Wiley & Sons Ltd, UK
- Wu, Y., Frizelle, G., Ayral, L., Marsein, J., Van de Merwe, E., Zhou, D. (2002), A Simulation Study on Supply Chain Complexity in Manufacturing Industry, In Proceedings of "Manufacturing Complexity Network Conference", 9-10 April 2002, Downing College, Cambridge, UK, Available online at:
http://www.ifm.eng.cam.ac.uk/mcn/pdf_files/part6_6.pdf (accessed 9 April 2007)
- Wu, Y., Frizelle, G., Efstathiou, J. (2007), A study on the cost of operational complexity in customer-supplier systems, *International Journal of Production Economics* 106 (1), 217-229
- Yang, B., Burns, N. (2003), Implications of postponement for the supply chain, *International Journal of Production Research* 41 (9), 2075-2090
- Ying, W., Dayong, S. (2005), Multi-agent framework for third party logistics in E-commerce, *Expert-Systems with Applications* 29 (2), 431-436
- Yu, S. B., Efstathiou J. (2002), An Introduction of Network Complexity, In Proceedings of "Manufacturing Complexity Network Conference", 9-10 April 2002, Downing College, Cambridge, UK, Available online at:
http://www.ifm.eng.cam.ac.uk/mcn/pdf_files/part6_7.pdf (accessed 9 April 2007)

Distance matrix from zone i to zone j (in km) for a case study in chapter 3.

Loc	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20	21	22	23	24	25	26	27	28	29	30	31	32	33	34	35	36
1	0.0	10.0	20.0	30.0	40.0	50.0	10.0	14.1	22.4	31.6	41.2	51.0	20.0	22.4	28.3	36.1	44.7	53.9	30.0	31.6	36.1	42.4	50.0	58.3	40.0	41.2	44.7	50.0	56.6	64.0	50.0	51.0	53.9	58.3	64.0	70.7
2	10.0	0.0	10.0	20.0	30.0	40.0	14.1	10.0	14.1	22.4	31.6	41.2	22.4	20.0	22.4	28.3	36.1	44.7	31.6	30.0	31.6	36.1	42.4	50.0	41.2	40.0	41.2	44.7	50.0	56.6	51.0	50.0	51.0	53.9	58.3	64.0
3	20.0	10.0	0.0	10.0	20.0	30.0	22.4	14.1	10.0	14.1	22.4	31.6	28.3	22.4	20.0	22.4	28.3	36.1	36.1	31.6	30.0	31.6	36.1	42.4	44.7	41.2	40.0	41.2	44.7	50.0	53.9	51.0	50.0	51.0	53.9	58.3
4	30.0	20.0	10.0	0.0	10.0	20.0	31.6	22.4	14.1	10.0	14.1	22.4	36.1	28.3	22.4	20.0	22.4	28.3	42.4	36.1	31.6	30.0	31.6	36.1	50.0	44.7	41.2	40.0	41.2	44.7	58.3	53.9	51.0	50.0	51.0	53.9
5	40.0	30.0	20.0	10.0	0.0	10.0	41.2	31.6	22.4	14.1	10.0	14.1	44.7	36.1	28.3	22.4	20.0	22.4	50.0	42.4	36.1	31.6	30.0	31.6	56.6	50.0	44.7	41.2	40.0	41.2	64.0	58.3	53.9	51.0	50.0	51.0
6	50.0	40.0	30.0	20.0	10.0	0.0	51.0	41.2	31.6	22.4	14.1	10.0	53.9	44.7	36.1	28.3	22.4	20.0	58.3	50.0	42.4	36.1	31.6	30.0	64.0	56.6	50.0	44.7	41.2	40.0	70.7	64.0	58.3	53.9	51.0	50.0
7	10.0	14.1	22.4	31.6	41.2	51.0	0.0	10.0	20.0	30.0	40.0	50.0	10.0	14.1	22.4	31.6	41.2	51.0	20.0	22.4	28.3	36.1	44.7	53.9	30.0	31.6	36.1	42.4	50.0	58.3	40.0	41.2	44.7	50.0	56.6	64.0
8	14.1	10.0	14.1	22.4	31.6	41.2	10.0	0.0	10.0	20.0	30.0	40.0	14.1	10.0	14.1	22.4	31.6	41.2	22.4	20.0	22.4	28.3	36.1	44.7	31.6	30.0	31.6	36.1	42.4	50.0	41.2	40.0	41.2	44.7	50.0	56.6
9	22.4	14.1	10.0	14.1	22.4	31.6	20.0	10.0	0.0	10.0	20.0	30.0	22.4	14.1	10.0	14.1	22.4	31.6	28.3	22.4	20.0	22.4	28.3	36.1	36.1	31.6	30.0	31.6	36.1	42.4	44.7	41.2	40.0	41.2	44.7	50.0
10	31.6	22.4	14.1	10.0	14.1	22.4	30.0	20.0	10.0	0.0	10.0	20.0	31.6	22.4	14.1	10.0	14.1	22.4	36.1	28.3	22.4	20.0	22.4	28.3	42.4	36.1	31.6	30.0	31.6	36.1	50.0	44.7	41.2	40.0	41.2	44.7
11	41.2	31.6	22.4	14.1	10.0	14.1	40.0	30.0	20.0	10.0	0.0	10.0	41.2	31.6	22.4	14.1	10.0	14.1	44.7	36.1	28.3	22.4	20.0	22.4	50.0	42.4	36.1	31.6	30.0	31.6	56.6	50.0	44.7	41.2	40.0	41.2
12	51.0	41.2	31.6	22.4	14.1	10.0	50.0	40.0	30.0	20.0	10.0	0.0	51.0	41.2	31.6	22.4	14.1	10.0	53.9	44.7	36.1	28.3	22.4	20.0	58.3	50.0	42.4	36.1	31.6	30.0	64.0	56.6	50.0	44.7	41.2	40.0
13	20.0	22.4	28.3	36.1	44.7	53.9	10.0	14.1	22.4	31.6	41.2	51.0	0.0	10.0	20.0	30.0	40.0	50.0	10.0	14.1	22.4	31.6	41.2	51.0	20.0	22.4	28.3	36.1	44.7	53.9	30.0	31.6	36.1	42.4	50.0	58.3
14	22.4	20.0	22.4	28.3	36.1	44.7	14.1	10.0	14.1	22.4	31.6	41.2	10.0	0.0	10.0	20.0	30.0	40.0	14.1	10.0	14.1	22.4	31.6	41.2	22.4	20.0	22.4	28.3	36.1	44.7	31.6	30.0	31.6	36.1	42.4	50.0
15	28.3	22.4	20.0	22.4	28.3	36.1	22.4	14.1	10.0	14.1	22.4	31.6	20.0	10.0	0.0	10.0	20.0	30.0	22.4	14.1	10.0	14.1	22.4	31.6	28.3	22.4	20.0	22.4	28.3	36.1	36.1	31.6	30.0	31.6	36.1	42.4
16	36.1	28.3	22.4	20.0	22.4	28.3	31.6	22.4	14.1	10.0	14.1	22.4	30.0	20.0	10.0	0.0	10.0	20.0	31.6	22.4	14.1	10.0	14.1	22.4	36.1	28.3	22.4	20.0	22.4	28.3	42.4	36.1	31.6	30.0	31.6	36.1
17	44.7	36.1	28.3	22.4	20.0	22.4	41.2	31.6	22.4	14.1	10.0	14.1	40.0	30.0	20.0	10.0	0.0	10.0	41.2	31.6	22.4	14.1	10.0	14.1	44.7	36.1	28.3	22.4	20.0	22.4	50.0	42.4	36.1	31.6	30.0	31.6
18	53.9	44.7	36.1	28.3	22.4	20.0	51.0	41.2	31.6	22.4	14.1	10.0	50.0	40.0	30.0	20.0	10.0	0.0	51.0	41.2	31.6	22.4	14.1	10.0	53.9	44.7	36.1	28.3	22.4	20.0	58.3	50.0	42.4	36.1	31.6	30.0
19	30.0	31.6	36.1	42.4	50.0	58.3	20.0	22.4	28.3	36.1	44.7	53.9	10.0	14.1	22.4	31.6	41.2	51.0	0.0	10.0	20.0	30.0	40.0	50.0	10.0	14.1	22.4	31.6	41.2	51.0	20.0	22.4	28.3	36.1	44.7	53.9
20	31.6	30.0	31.6	36.1	42.4	50.0	22.4	20.0	22.4	28.3	36.1	44.7	14.1	10.0	14.1	22.4	31.6	41.2	10.0	0.0	10.0	20.0	30.0	40.0	50.0	14.1	10.0	14.1	22.4	31.6	41.2	22.4	20.0	22.4	28.3	36.1
21	36.1	31.6	30.0	31.6	36.1	42.4	28.3	22.4	20.0	22.4	28.3	36.1	22.4	14.1	10.0	14.1	22.4	31.6	20.0	10.0	0.0	10.0	20.0	30.0	40.0	22.4	14.1	10.0	14.1	22.4	31.6	28.3	22.4	20.0	22.4	28.3
22	42.4	36.1	31.6	30.0	31.6	36.1	36.1	28.3	22.4	20.0	22.4	28.3	31.6	22.4	14.1	10.0	14.1	22.4	30.0	20.0	10.0	0.0	10.0	20.0	31.6	22.4	14.1	10.0	14.1	22.4	36.1	28.3	22.4	20.0	22.4	28.3
23	50.0	42.4	36.1	31.6	30.0	31.6	44.7	36.1	28.3	22.4	20.0	22.4	41.2	31.6	22.4	14.1	10.0	14.1	40.0	30.0	20.0	10.0	0.0	10.0	41.2	31.6	22.4	14.1	10.0	14.1	44.7	36.1	28.3	22.4	20.0	22.4
24	58.3	50.0	42.4	36.1	31.6	30.0	53.9	44.7	36.1	28.3	22.4	20.0	51.0	41.2	31.6	22.4	14.1	10.0	50.0	40.0	30.0	20.0	10.0	0.0	51.0	41.2	31.6	22.4	14.1	10.0	53.9	44.7	36.1	28.3	22.4	20.0
25	40.0	41.2	44.7	50.0	56.6	64.0	30.0	31.6	36.1	42.4	50.0	58.3	20.0	22.4	28.3	36.1	44.7	53.9	10.0	14.1	22.4	31.6	41.2	51.0	0.0	10.0	20.0	30.0	40.0	50.0	10.0	14.1	22.4	31.6	41.2	51.0
26	41.2	40.0	41.2	44.7	50.0	56.6	31.6	30.0	31.6	36.1	42.4	50.0	22.4	20.0	22.4	28.3	36.1	44.7	14.1	10.0	14.1	22.4	31.6	41.2	10.0	0.0	10.0	20.0	30.0	40.0	14.1	10.0	14.1	22.4	31.6	41.2
27	44.7	41.2	40.0	41.2	44.7	50.0	36.1	31.6	30.0	31.6	36.1	42.4	28.3	22.4	20.0	22.4	28.3	36.1	22.4	14.1	10.0	14.1	22.4	31.6	20.0	10.0	0.0	10.0	20.0	30.0	22.4	14.1	10.0	14.1	22.4	31.6
28	50.0	44.7	41.2	40.0	41.2	44.7	42.4	36.1	31.6	30.0	31.6	36.1	36.1	28.3	22.4	20.0	22.4	28.3	31.6	22.4	14.1	10.0	14.1	22.4	30.0	20.0	10.0	0.0	10.0	20.0	31.6	22.4	14.1	10.0	14.1	22.4
29	56.6	50.0	44.7	41.2	40.0	41.2	50.0	42.4	36.1	31.6	30.0	31.6	44.7	36.1	28.3	22.4	20.0	22.4	41.2	31.6	22.4	14.1	10.0	14.1	40.0	30.0	20.0	10.0	0.0	10.0	41.2	31.6	22.4	14.1	10.0	14.1
30	64.0	56.6	50.0	44.7	41.2	40.0	58.3	50.0	42.4	36.1	31.6	30.0	53.9	44.7	36.1	28.3	22.4	20.0	51.0	41.2	31.6	22.4	14.1	10.0	50.0	40.0	30.0	20.0	10.0	0.0	51.0	41.2	31.6	22.4	14.1	10.0
31	50.0	51.0	53.9	58.3	64.0	70.7	40.0	41.2	44.7	50.0	56.6	64.0	30.0	31.6	36.1	42.4	50.0	58.3	20.0	22.4	28.3	36.1	44.7	53.9	10.0	14.1	22.4	31.6	41.2	51.0	0.0	10.0	20.0	30.0	40.0	50.0
32	51.0	50.0	51.0	53.9	58.3	64.0	41.2	40.0	41.2	44.7	50.0	56.6	31.6	30.0	31.6	36.1	42.4	50.0	22.4	20.0	22.4	28.3	36.1	44.7	14.1	10.0	14.1	22.4	31.6	41.2	10.0	0.0	10.0	20.0	30.0	40.0
33	53.9	51.0	50.0	51.0	53.9	58.3	44.7	41.2	40.0	41.2	44.7	50.0	36.1	31.6	30.0	31.6	36.1	42.4	28.3	22.4	20.0	22.4	28.3	36.1	22.4	14.1	10.0	14.1	22.4	31.6	20.0	10.0	0.0	10.0	20.0	30.0
34	58.3	53.9	51.0	50.0	51.0	53.9	50.0	44.7	41.2	40.0	41.2	44.7	42.4	36.1	31.6	30.0	31.6	36.1	28.3	22.4	20.0	22.4	28.3	31.6	22.4	14.1	10.0	14.1	22.4	30.0	20.0	10.0	0.0	10.0	20.0	
35	64.0	58.3	53.9	51.0	50.0	51.0	56.6	50.0	44.7	41.2	40.0	41.2	50.0	42.4	36.1	31.6	30.0	31.6	44.7	36.1	28.3	22.4	20.0	22.4	41.2	31.6	22.4	14.1	10.0	14.1	40.0	30.0	20.0	10.0	0.0	
36	70.7	64.0	58.3	53.9	51.0	50.0	64.0	56.6	50.0	44.7	41.2	40.0	58.3	50.0	42.4	36.1	31.6	30.0	53.9	44.7	36.1	28.3	22.4	20.0	51.0	41.2	31.6	22.4	14.1	10.0	50.0	40.0	30.0	20.0	10.0	0.0

Derivation of the Equations for the Slope and Intercept for Regression Analysis

In this appendix, we excerpt the content from Nahmias (2005), pp. 104-106, to demonstrate the method for deriving the optimal values for the slope and intercept of the equations in regression analysis that we use in chapter 6. Assume that the data are $(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)$, and the regression model to be fitted in $Y = a + bX$. Define

$$g(a, b) = \sum_{i=1}^n [y_i - (a + bx_i)]^2$$

Interpret $g(a, b)$ as the sum of the squares of the distance from the line $a + bx$ to the data point points y_i . The object of the analysis is to choose a and b to minimise $g(a, b)$. This is accomplished where

$$\frac{\partial g}{\partial a} = \frac{\partial g}{\partial b} = 0$$

That is,

$$\frac{\partial g}{\partial a} = -\sum_{i=1}^n 2[y_i - (a + bx_i)] = 0$$

$$\frac{\partial g}{\partial b} = -\sum_{i=1}^n 2x_i[y_i - (a + bx_i)] = 0$$

which results in the two equations

$$an + b \sum_{i=1}^n x_i = \sum_{i=1}^n y_i \quad (\text{AB.1})$$

$$a \sum_{i=1}^n x_i + b \sum_{i=1}^n x_i^2 = \sum_{i=1}^n x_i y_i \quad (\text{AB.2})$$

These are two linear equations in the unknowns a and b . Multiplying Equation (AB.1) by $\sum x_i$ and Equation (AB.2) by n gives

$$an \sum_{i=1}^n x_i + b \left(\sum_{i=1}^n x_i \right)^2 = \left(\sum_{i=1}^n x_i \right) \left(\sum_{i=1}^n y_i \right) \quad (\text{AB.3})$$

$$an \sum_{i=1}^n x_i + bn \sum_{i=1}^n x_i^2 = n \sum_{i=1}^n x_i y_i \quad (\text{AB.4})$$

Subtracting Equation (AB.3) from Equation (AB.4) results in

$$b \left[n \sum_{i=1}^n x_i^2 - \left(\sum_{i=1}^n x_i \right)^2 \right] = n \sum_{i=1}^n x_i y_i - \left(\sum_{i=1}^n x_i \right) \left(\sum_{i=1}^n y_i \right) \quad (\text{AB.5})$$

Define $S_{xy} = n \sum x_i y_i - (\sum x_i)(\sum y_i)$ and $S_{xx} = n \sum x_i^2 - (\sum x_i)^2$. It follows that Equation (AB.5) may be written $bS_{xx} = S_{xy}$, which gives

$$b = \frac{S_{xy}}{S_{xx}} \quad (\text{AB.6})$$

From Equation (AB.1) we have

$$ay = \sum_{i=1}^n y_i - b \sum_{i=1}^n x_i$$

or

$$a = \bar{y} - b\bar{x} \quad (\text{AB.7})$$

where $\bar{y} = (1/n) \sum y_i$ and $\bar{x} = (1/n) \sum x_i$.

These formulas can be specialised to the forecasting problem when the independent variable is assumed to be time. In that case, the data are of the form $(1, D_1), (2, D_2), \dots, (n, D_n)$, and the forecasting equation is of the form $\hat{D}_t = a + bt$. The various formulas can be simplified as follows:

$$\begin{aligned} \sum x_i &= 1 + 2 + 3 + \dots + n = \frac{n(n+1)}{2} \\ \sum x_i^2 &= 1 + 4 + 9 + \dots + n^2 = \frac{n(n+1)(2n+1)}{6} \end{aligned}$$

Hence, we can write

$$\begin{aligned} S_{xy} &= n \sum_{i=1}^n i D_i - \frac{n(n+1)}{2} \sum_{i=1}^n D_i \\ S_{xx} &= \frac{n^2(n+1)(2n+1)}{6} - \frac{n^2(n+1)^2}{4} \\ b &= \frac{S_{xy}}{S_{xx}} \\ a &= \bar{D} - \frac{b(n+1)}{2} \end{aligned}$$