

ECONOMETRIC MODELLING OF HOUSEHOLD BEHAVIOUR

by

AUDRIENNE CUTAJAR BEZZINA

Department of Mathematics

University of Malta

September 1999

A Dissertation submitted to the

Faculty of Science

in partial fulfilment of the

requirements for the degree of

Masters of Science in Mathematics



University of Malta Library – Electronic Thesis & Dissertations (ETD) Repository

The copyright of this thesis/dissertation belongs to the author. The author's rights in respect of this work are as defined by the Copyright Act (Chapter 415) of the Laws of Malta or as modified by any successive legislation.

Users may access this full-text thesis/dissertation and can make use of the information contained in accordance with the Copyright Act provided that the author must be properly acknowledged. Further distribution or reproduction in any format is prohibited without the prior permission of the copyright holder.

TO MY PARENTS WHO WERE MY FIRST EDUCATORS

ABSTRACT

The aim of this dissertation is to first study the theory related to consumer expenditure and then try and fit the best statistical distribution through the data points obtained from the survey.

The aim of chapter 1 is to justify the existence of rules for deciding how much a consumer spends on a particular good. These rules will be defined as demand functions. The set of preferences of a consumer will be defined, and hence, one can see how from these preferences, a utility function is obtained. The utility function models how a household spends its budget. Optimisation of these utility functions will yield demand equations. The concept of duality will be then introduced and hence the transition from one set of demands to the other will be explained. Elasticities and their use will also be explained in this chapter. Some constraints of the demand equations will be reviewed. These constraints are the basis for the models mentioned in chapter 2.

The theory presented in chapter 2 offers a framework within which one can attempt to organise and interpret the data. In chapter 2, I shall discuss the application of some models, each of which seeks to describe the way in which income is spent. These models not only show how the theory can be made specific, but also allow us to check the usefulness of the theory in empirical work. The models will be presented in rough chronological order. The first part of the chapter is devoted to the 'Linear

Expenditure System' (LES), the first practical model to be based entirely upon the theory. Since this model is the heart of all the models (mainly because all new models are almost always derived from the LES), the full theory of the LES shall be reproduced in this chapter. Then a review of other models derived from the LES, namely the 'Rotterdam Model', the 'Transcendental Logarithmic Model', the 'Almost Ideal Demand System', and 'An implicitly additive Demand System' is carried out. For the last four models, the authors estimated a general functional form with and without the restrictions, thus testing whether the restrictions were validly imposed. In that section, I shall first state the functional form used by the authors, then review any results they obtained from their testing. In the last part of the chapter I shall review a counter argument to the use of the above models. In fact some authors write that there is no need to specify a model of micro-economic behaviour because some of the coefficients for the macro-economic relation can be estimated using the cross-section data.

Although there exists quite a number of demand functions, I have found out that no applied study was done on the Household Budgetary Survey for Malta. A problem coming from economics, that can be solved statistically is a problem associated with the data. The data obtained from surveys is considered to be a non-realistic data, mainly because people tend to exaggerate their expenditure, and minimise their income. To solve the above problem, one needs to look into density estimation. In the first part of chapter 3, current methods being used for density estimation will be reviewed, with some of these methods being tried out on Maltese data. In the second part of the chapter, I shall try to fit a curve by using simple programs written by

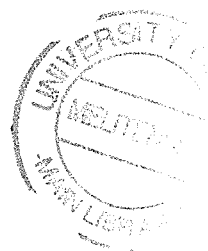
myself in Qbasic©. A curve is fitted for each and every commodity, and a detailed study of the outliers is given after each commodity.

ACKNOWLEDGMENTS

I would like to thank my husband, Reuben for his patience, his help and for having believed in me even when I didn't believe in myself.

Particular mention also goes to Mr. Alfred Camilleri for making the data available, Dr. Joseph Muscat, M.A. (Hons) (Oxon), Ph.D. (Princeton), Mr. James Borg, B.Sc., M.Sc., Prof. Edward Scicluna and Mr. Etienne Caruana, B.Sc. for the help offered throughout the study. Last but not least I wish to thank Prof. Stanley Fiorini, B.Phil (Ileythrop), M.A.(Oxon),Ph.D. (Open) for making the course available.

I am also indebted to all other persons who in any way contributed towards the completion of this dissertation.



Contents

CONTENTS

CHAPTER 1	FROM UTILITY TO DEMAND	1
1.1	Preference Relations	2
1.2	Utility Functions	6
1.3	Demand Functions	8
1.4	Cost Minimisation and the Cost Function	9
1.5	Properties of the Cost Function	10
1.6	Justification for Utility Maximisation	13
1.7	Elasticities	19
1.8	Constraints for Demand Functions	20
CHAPTER 2	MODELLING HOUSEHOLD BEHAVIOUR	25
2.1	Models for Demand Equations	25
2.1.1	The Linear Expenditure System	26
2.1.1.1	Testing the Restrictions	27
2.1.2	The Rotterdam Model	37
2.1.2.1	Testing the Restrictions	38
2.2	Flexible Functional Forms	41
2.2.1	The Transcendental Logarithmic Model	42
2.2.1.1	Testing the Restrictions	43

2.2.2	The Almost Ideal Demand System	48
2.2.2.1	Testing the Restrictions	52
2.2.3	Conclusions for the Above Three Models	53
2.2.4	An Implicitly Additive Demand System	54
2.2.4.1	Testing the Restrictions	57
2.3	Conclusion	58
CHAPTER 3	DENSITY ESTIMATION AND CURVE FITTING	60
3.1	Basic Definitions and Properties	61
3.2	Measurements of Error	63
3.3	Non-Parametric Estimation	64
3.3.1	The Histogram	64
3.3.2	The Naïve Estimator	68
3.3.3	The Kernel Estimator	70
3.4	Parametric Density Estimation	73
3.4.1	The Least Squares Method	73
3.5	Curve Fitting Using Parametric Estimation	74
3.5.1	Expenditure	75
3.5.2	Food	77
3.5.3	Beverages and Tobacco	80
3.5.4	Clothing and Footwear	82
3.5.5	Housing	85
3.5.6	Fuel and Power	87
3.5.7	Furniture and Household Equipment	90
3.5.8	Transport and Communication	92

3.5.9	Personal Care and Health	95
3.5.10	Education and Recreation	97
3.5.11	Other Goods and Services	100
3.6	Improvements	102
REFERENCES		105
SOFTWARE USED FOR DATA		109
ANALYSIS		
APPENDIX A		111
A.1	Program for Naïve Estimation written using Turbo Pascal Version 7	111
A.2	Curve Fitting Program for Total Expenditure Written Using Quick Basic	113
A.3	Curve Fitting Program for Food Budget Share Written Using Quick Basic	115
A.4	Curve Fitting Program for Beverages and Tobacco Budget Share Written Using Quick Basic	117
A.5	Curve Fitting Program for Clothing and Footwear Budget Share Written Using Quick Basic	119
A.6	Curve Fitting Program for Housing Budget Share Written Using Quick Basic	120

A.7	Curve Fitting Program for Fuel and Power Budget Share Written Using Quick Basic	121
A.8	Curve Fitting Program for Furniture and Household Equipment Budget Share Written Using Quick Basic	123
A.9	Curve Fitting Program for Transport and Communication Budget Share Written Using Quick Basic	125
A.10	Curve Fitting Program for Personal Care and Health Budget Share Written Using Quick Basic	127
A.11	Curve Fitting Program for Education and Recreation Budget Share Written Using Quick Basic	129
A.12	Curve Fitting Program for Other Goods and Services Budget Share Written Using Quick Basic	130

Chapter 1

FROM UTILITY TO DEMAND

Chapter 1

FROM UTILITY TO DEMAND

The aim of this chapter is to justify the existence of rules for deciding how much a consumer spends on a particular good. These rules will be defined as demand functions. I shall start by defining the preferences of a consumer. From these preferences, one can obtain a utility function, which is the function that models how a household spends its budget. Optimisation of these utility functions will yield demand equations.

The concept of duality will be then introduced and hence the transition from one set of demands to the other will be explained. In the last part of the chapter, elasticities and some constraints of the demand equations will be reviewed. These constraints are the basis for the models mentioned in chapter 2.

1.1 Preference Relations

Why does a household prefer a particular brand of food? Is it for the price or the quality? Why does the household spend money on clothes instead of spending them on food? In order to answer these questions, one needs to define some sort of function that is unique to a particular household, and that it identifies the particular choices that a household makes. This function is called the **utility function**. In order to get demand equations from a utility function, one needs to maximise the utility subject to a constraint. In the first part of this chapter, the constraint used will be a linear constraint, namely that

$$y = \sum_{i=1}^n p_i x_i$$

where y is the total expenditure

p_i 's are the prices to buy the quantities x_i 's.

So, the aim of the first part of this chapter is to review how various authors answered the question that 'Given a household with a budget or total expenditure of y , which is to be spent in a given period of time on some or on all of l goods, what will be the household's preferences? The following diagram shows an example on two commodities.

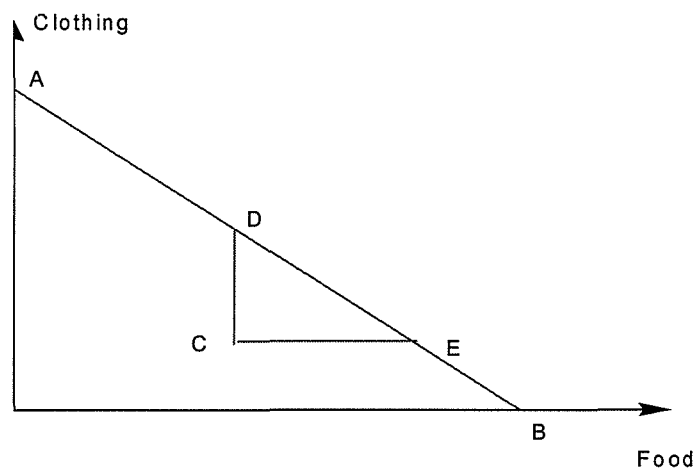


Figure 1.1 *A linear budget constraint with a survival constraint*

At A, one can easily see that all the budget is spent on clothes, while at B, all the budget is spent on food. However, if a restriction is made such that a household to survive needs minimum quantities of the two goods, (in this case at C), then choice will be restricted to triangle CDE.

The transition from preference relations to demand functions is best shown using the following diagram.

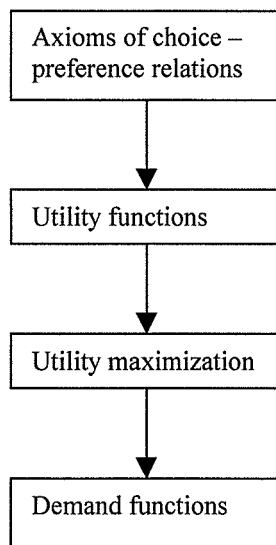


Figure 1.2 *The transition from preference to demand functions*

Up to now, preference is just a word with no mathematical explanation. So, define a finite number l of commodities useful for the consumer such as food, clothing, etc. Thus, a commodity bundle will be a sequence of real numbers $x_1, x_2, \dots, x_l$ indicating the quantity of each commodity. The price of a commodity $p_1, p_2, \dots, p_h$ is a real number indicating the amount paid for one unit of the commodity. Then the value of a commodity bundle given a price vector p is

$$\sum_{h=1}^l p_h x_h = p \cdot x$$

The set of all possible consumption bundles is defined as X - the consumption set, where X is closed, convex and bounded below.

A set X is said to be convex if a line segment between any two points x_1 and x_2 belongs entirely to that set. In terms of commodity bundles, if there are two

commodities x and y that are indifferent to each other, then a linear combination of these two (or $\lambda x + (1 - \lambda)y$) is preferred to x and y .

The preference relation $\geq$ will be defined as a binary relation on $(X \times X)$ by

$$\forall x, y: x \in X, y \in X, x \geq y \Leftrightarrow x \text{ is at least as good as } y$$

Having defined a preference ordering, some axioms are needed to transform this ordering into the utility function. The following four axioms are sufficient axioms that allow the preference order to be represented by a utility function. These axioms are known as the axioms of preferences.

Axiom 1: Reflexivity - $\forall x \in X, x \geq x$

i.e. each bundle is as good as itself

Axiom 2: Completeness - $\forall x, y \in X, x \geq y \cup y \geq x$

i.e. that any two bundles can be compared, and the consumer can judge between any two bundles.

Axiom 3: Transitivity or consistency - $\forall x, y, z \in X: x \geq y, y \geq z \Rightarrow x \geq z$

i.e. If the consumer prefers A to B, but prefers C to A, then one would expect that he prefers C to B.

Axiom 4: Continuity - $\left\{ \forall x \in X, \left\{ y \in X \mid y \geq x \right\}, \left\{ y \in X \mid x \geq y \right\} \right\}$ are closed relative to X .

i.e. the specified sets contain their own boundaries.

Bundle x is strictly preferred to bundle y , written as $x > y$ if $x \geq y$ and not $y \geq x$.

Bundle x is said to be indifferent to bundle y , written as $x \sim y$ if $x \geq y$ and $y \geq x$.

Then $\sim$ defines an equivalence relation on X since it is reflexive, symmetric and transitive.

Trivially, it is reflexive since $x \sim x$ if $x \geq x$ or $x \geq x$.

If it were symmetric then $x \sim y$ would imply $y \sim x$. But $x \sim y$ means that $x \geq y$ and $y \geq x$ which is the same as saying that $y \geq x$ and $x \geq y$ which imply that $y \sim x$.

If it were transitive, then if $x \sim y$ and $y \sim z$, this would imply that $x \sim z$. But $x \sim y$ means that $x \geq y$ and $y \geq x$, and $y \sim z$ means that $y \geq z$ and $z \geq y$, and because the preference ordering was defined to be transitive, then the indifference relation is transitive as well.

The previous four axioms are enough to represent the preference ordering by a utility function, however, the function defined is not unique. To guarantee that the best choice will lie on, rather than inside the budget constraint, another axiom is defined.

Axiom 5: Nonsatiation – The utility function is non-decreasing in each of its arguments, and is increasing for at least one of its arguments.

1.2 Utility Functions

In order to go from preferences to utility functions, we need to make use of the following representation theorem.

Theorem: Let X be a completely ordered, separable and connected space. If for every $x' \in X$ the sets $\{x \in X | x \leq x'\}$ and $\{x \in X | x' \leq x\}$ are closed, there exists on X a continuous, real, order-preserving function.

If X denotes a set and $\geq$ a binary relation on X , then a function u from X into the real line $\mathfrak{R}$ is a **utility function** if for any two points x and y , $u(x) \geq u(y)$ if and only if $x \geq y$. Perhaps, a more appropriate name for utility functions would be ‘preference representation functions’. These functions are ordinal and not cardinal functions, since only their position in the ordering is important, and not the actual value of the function. Besides, for any utility function u and any increasing transformation $f: \mathfrak{R} \longrightarrow \mathfrak{R}$, the function $v = f \circ u$ will also be a utility function for the same relation $\geq$.

In the next section, the utility functions will be maximised to obtain the demand equations. The demand equation is defined as a function of the price p and income y , namely

$$x_i = g_i(y, p)$$

To ensure that the maximum of the above equation is in fact a global maximum, one needs that the set on which the function is defined is in fact a convex set. This will justify the following axiom, which is added to the five axioms mentioned in the previous section.

Axiom 6: Strict Convexity – If two bundles x and y are indifferent, then a linear combination of these bundles is preferred to x and y .

This means that preferences are convex if

$$\lambda x + (1 - \lambda)y > x \text{ when } x \sim y$$

But the fact that

$$u(x) > u(y) \text{ when } x > y$$

implies that

$$u(\lambda x + (1 - \lambda)y) > u(y) \text{ when } x \sim y$$

.....(1.1)

A utility function that has property (1.1) is said to be *strictly quasi-concave*.

Generally speaking, a function $u: X \longrightarrow \mathbb{R}$ is called *quasi-concave* if $u(\lambda x + (1 - \lambda)y) \geq \min\{u(x), u(y)\}$ for all $x, y \in X$ and any $\lambda, 0 \leq \lambda \leq 1$.

1.3

Demand Functions

As mentioned earlier, the demand equation is a function of the price of the commodity and the income of the consumer buying that commodity. To obtain the demand equations, one has to maximise the utility functions subject to a constraint. The constraint considered will be the linear constraint. So the problem will reduce to maximisation of $u(x)$ subject to $\sum p_k x_k = y$. The case considered here is one where the budget line is tangent to a smoothly convex indifference curve at a single point, both goods are bought, demand is continuous, and can be solved by simple differential calculus.

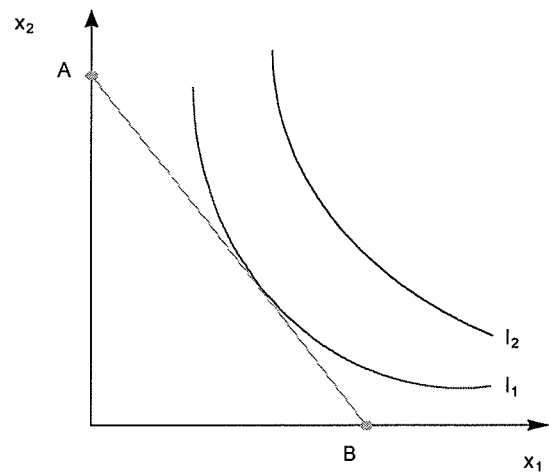


Figure 1.3

Indifference curve with a tangent budget line

The solution of the maximisation problem will give rise to a system of demand equations, referred to as the Marshallian demand functions defined by

$$x_i = g_i(y, p)$$

Using Lagrange multipliers, the equation is transformed to

$$\phi = u(x) + \lambda(y - \sum p_k x_k)$$

and on partial differentiation

$$\frac{\partial u(x)}{\partial x_i} = \lambda p_i \quad \text{and} \quad y = \sum p_k x_k$$

where λ is the marginal utility of total expenditure; λ is the amount by which the maximand would increase given a unit increase in y .

1.4 *Cost Minimisation and the Cost Function*

Suppose now that the problem is expressed in another way, namely that the consumer must select goods in such a way as to minimise the outlay necessary to reach the utility u . This will give rise to the same vector of commodities as in the previous section. The problem is therefore to minimise $y = \sum p_k x_k$ subject to $u(x) = u$. This is also referred to as the dual problem. The same solution is obtained, but whereas in the Marshallian demand functions, the demands are a function of y and p , in these new demand functions, also called as the Hicksian demand functions, the demand will be a function of u and p , i.e. $x_i = h_i(u, p)$.

On substituting back,

$$u = u(x_1, x_2, \dots, x_n) = u[g_1(y, p), g_2(y, p), \dots, g_n(y, p)] = \phi(y, p)$$

$$y = \sum p_k h_k(u, p) = c(u, p)$$

where

1. $\varphi(y, p)$ is maximum attainable utility called - *indirect utility function*.

$$\varphi(y, p) = \max_x [\nu(x); \sum p_k x_k = y]$$

2. $c(u, p)$ is the minimum cost of attaining u at prices p - *cost function*.

$$c(u, p) = \min_x [\sum p_k x_k; \nu(x) = u]$$

So, basically, the indirect utility function and the cost function are related as shown in the following diagram.

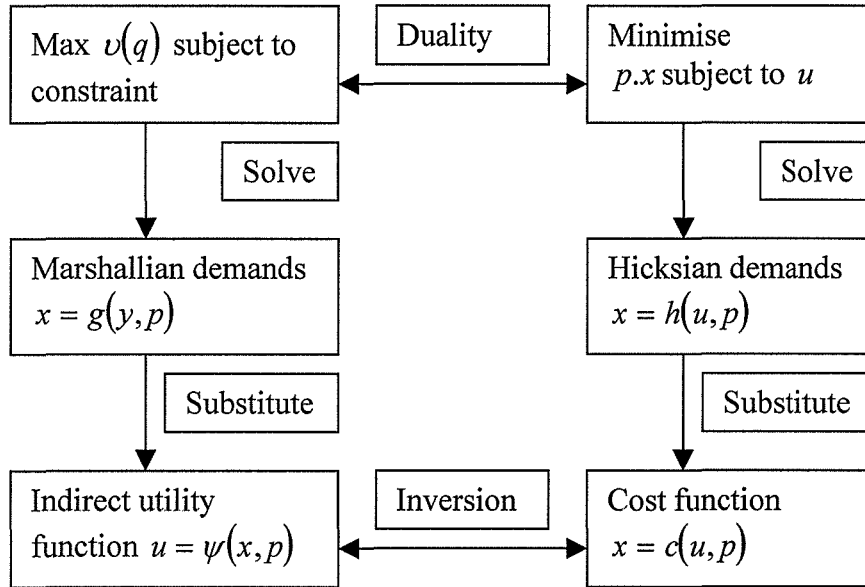


Figure 1.4 Utility maximisation and cost minimisation⁹

1.5 Properties of the Cost Function

1. The cost function is homogeneous of degree one for a scalar $\theta > 0$, i.e.

$c(u, \theta p) = \theta c(u, p)$. This means that if the prices double, twice as much outlay is necessary to stay on the same indifference curve.

Using the fact that

$$c(u, p) = \sum p_k h_k(u, p)$$

then

$$c(u, \theta p) = \sum \theta p_k h_k(u, \theta p)$$

and since θ is a scalar greater than zero, then

$$\sum \theta p_k h_k(u, \theta p) = \theta \sum p_k h_k(u, p) = \theta c(u, p)$$

2. The cost function is increasing in u , non-decreasing in p and increasing in at least one price.

Remember that the Nonsatiation axiom read that the utility function is non-decreasing in each of its arguments, and is increasing for at least one of its arguments. Therefore, since $c(u, p) = \sum p_k h_k(u, p)$, then the above property follows immediately.

3. The cost function is concave in prices. This implies that as prices rise, the cost rises no more than linearly.

Consider two price vectors p^1 and p^2 and let $p^3 = \lambda p^1 + (1 - \lambda)p^2$. Let x^3 be the optimal choice of commodities when prices are p^3 and utility is u . Then by

$$y = \sum p_k h_k(u, p) = c(u, p),$$

$$c(u, p^3) = \sum x_k^3 p_k^3 = \lambda \sum x_k^3 p_k^1 + (1 - \lambda) \sum x_k^3 p_k^2$$

But x^3 is not necessarily the cost minimising bundle for u at either p^1 or p^2 , such that

$$c(u, p^1) \leq \sum x_k^3 p_k^1$$

and

$$c(u, p^2) \leq \sum x_k^3 p_k^2.$$

Hence

$$c(u, \lambda p^1 + (1 - \lambda)p^2) \geq \lambda c(u, p^1) + (1 - \lambda)c(u, p^2)$$

which establishes concavity.

4. Where they exist, the partial derivatives of the cost function with respect to prices are the Hicksian demand functions, i.e.

$$\frac{\partial c(u, p)}{\partial p_i} \equiv h_i(u, p) = x_i$$

Consider an arbitrary vector of prices p^1 , a utility level u , and the corresponding vector of optimal choices x^1 . For any other price p , define the function $Z(p)$ by

$$Z(p) = \sum_k p_k x_k^1 - c(u, p)$$

Since x^1 is not necessarily optimal for p , the cost of x^1 at p must always be at least as great as the cost of the optimal vector p . Hence $Z(p)$ is always greater or equal to zero, or it attains its minimum when $p = p^1$. Hence, if the derivative exists at p^1 ,

$$\frac{\partial Z(p^1)}{\partial p_i} = x_i^1 - \frac{\partial c}{\partial p_i}(u, p^1) = 0$$

and since p^1 is arbitrary, shows that the partial derivatives are the Hicksian demand functions.

In Figure 1.4, one can see that the cost and the indirect utility function are inverses of each other. Thus if we are given an indirect utility function, it is possible to obtain the Marshallian demands by making use of the following steps.

The cost and indirect utility functions are inverses thus

$$\phi[c(u, p), p] \equiv u$$

On differentiating the function with respect to p_i and keeping u constant,

$$\frac{\partial \phi}{\partial y} \cdot \frac{\partial c}{\partial p_i} + \frac{\partial \phi}{\partial p_i} = 0$$

On re-arranging,

$$x_i = g_i(y, p) = -\frac{\partial \phi / \partial p_i}{\partial \phi / \partial y}$$

This is called the Roy's identity.

Figure 1.5 explains how the demand, the cost and the indirect utility functions are related.

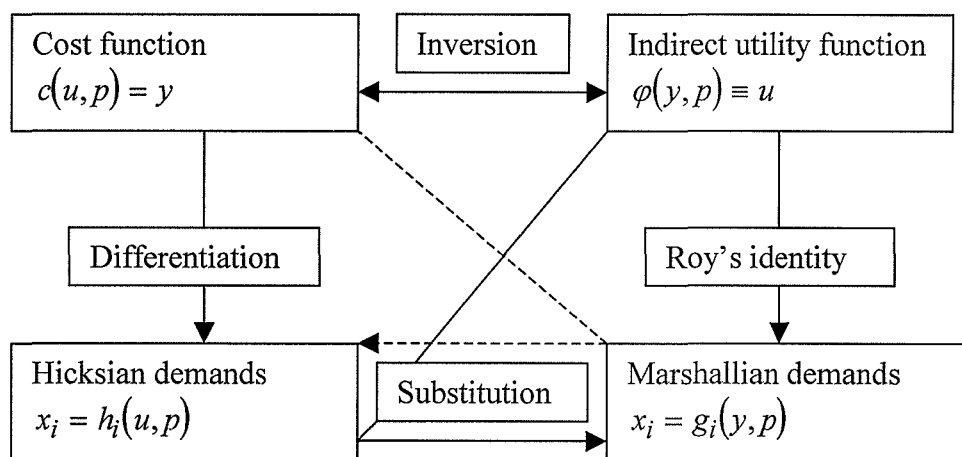


Figure 1.5 Changing from demand, cost and indirect utility functions.⁹

1.6 Justification for Utility Maximisation

The entire preceding discussion has been based on the assumption that the second order conditions for a local maximum and the conditions for a global maximum are

satisfied. In fact, an important result by Philips³¹ regarding sufficient conditions for the second order differential equations to be a maximum is:

'for a constrained maximum, with one constraint, the determinant of the bordered Hessian should have the sign of $(-1)^n$, where n is the number of variables, the largest principal minor should have a sign opposite to this, and successively smaller principal minors should alternate in sign, down to the principal minor of order 2'

A Hessian matrix is defined to be a matrix of second-order partial derivatives of the utility function, i.e.

$$\frac{\partial^2 u}{\partial x_i \partial x_j}$$

By Young's Theorem

$$\frac{\partial^2 u}{\partial x_i \partial x_j} = \frac{\partial^2 u}{\partial x_j \partial x_i}$$

this will imply that a Hessian matrix is symmetric.

For ease of notation, define

$$u_{ii} = \frac{\partial^2 u}{\partial x_i^2} \text{ and } u_{ij} = \frac{\partial^2 u}{\partial x_i \partial x_j}$$

Thus a Hessian matrix is defined as

$$U = \begin{bmatrix} u_{11} & u_{12} & \cdot & \cdot & u_{1n} \\ u_{21} & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot \\ u_{n1} & \cdot & \cdot & \cdot & u_{nn} \end{bmatrix}$$

A bordered Hessian matrix is a Hessian matrix bordered by a row and a column containing a zero and the partial derivatives (with respect to x_i) of the budget constraint $\sum_i p_i x_i - y = 0$.

Since $\frac{\partial(\sum_i p_i x_i - y)}{\partial x_i} = p_i$, then the bordered Hessian matrix can be defined as

$$\begin{bmatrix} 0 & p_1 & p_2 & \cdot & \cdot & p_n \\ p_1 & u_{11} & \cdot & \cdot & \cdot & u_{1n} \\ p_2 & u_{21} & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ p_n & u_{n1} & \cdot & \cdot & \cdot & u_{nn} \end{bmatrix}$$

Consider the case when we have two variables, i.e. the determinant of the bordered Hessian matrix will be positive $(-1)^2$ if the equations obtained are in fact maximum.

Thus the sufficient conditions will be that:

$$\begin{vmatrix} 0 & p_1 & p_2 \\ p_1 & u_{11} & u_{12} \\ p_2 & u_{21} & u_{22} \end{vmatrix} = \{-p_1(p_1 u_{22} - p_2 u_{12}) + p_2(p_1 u_{21} - p_2 u_{11})\} > 0$$

$$-p_1^2 u_{22} + p_1 p_2 u_{12} + p_1 p_2 u_{21} - p_2^2 u_{11} > 0$$

$$2p_1 p_2 u_{12} - p_1^2 u_{22} - p_2^2 u_{11} > 0$$

and that

$$\begin{vmatrix} 0 & p_1 \\ p_1 & u_{11} \end{vmatrix} = -p_1^2 < 0$$

But the determinant of the second principal minor is surely negative, which agrees with the second part of the stated theorem. Thus all there is to prove is that:

$$2p_1p_2u_{12} - p_1^2u_{22} - p_2^2u_{11} > 0 \quad \dots\dots\dots(1.2)$$

Now, a matrix A is defined to be negative semi-definite if it satisfies the condition $z'Az \leq 0$ for any vector $z \neq 0$.

Eggleston¹³ in his book proved that if a utility function is strictly concave, then the Hessian matrix is semi-definite. If this were the case, then writing the Hessian matrix

U , and the vector $\begin{bmatrix} p_1 \\ p_2 \end{bmatrix}$ in the following way, one can see that (1.2) is proved.

For

$$\begin{bmatrix} p_1 & p_2 \end{bmatrix} \begin{bmatrix} u_{11} & u_{12} \\ u_{21} & u_{22} \end{bmatrix} \begin{bmatrix} p_1 \\ p_2 \end{bmatrix} = \begin{bmatrix} p_1u_{11} + p_2u_{21} & p_1u_{12} + p_2u_{22} \end{bmatrix} \begin{bmatrix} p_1 \\ p_2 \end{bmatrix}$$

$$p_1^2u_{11} + p_1p_2u_{21} + p_1p_2u_{12} + p_2^2u_{22} \leq 0$$

$$p_1^2u_{11} + p_2^2u_{22} \leq -2p_1p_2u_{21}$$

Therefore

$$p_1^2u_{11} + p_2^2u_{22} - 2p_1p_2u_{21} \leq -2p_1p_2u_{21} - 2p_1p_2u_{21} \leq -4p_1p_2u_{21} < 0.$$

where p_1, p_2 are always positive.

So, the sufficient condition would be proved for a strictly concave utility function.

However, concavity is too restrictive for theoretical purposes, so it is more general to assume that the utility function is quasi-concave. The sufficient conditions are also satisfied for quasi concavity.

Since results are being proved for two variables ($n=2$), consider the utility function as a function of x_1 and x_2 , i.e. $u = f(x_1, x_2)$.

Then

$$du = \frac{\partial u}{\partial x_1} dx_1 + \frac{\partial u}{\partial x_2} dx_2 = 0$$

$$u_1 dx_1 + u_2 dx_2 = 0$$

$$\frac{dx_2}{dx_1} = -\frac{u_1}{u_2}$$

If the utility function were convex, then

$$\frac{d^2 x_2}{dx_1^2} > 0$$

But

$$\frac{d^2 x_2}{dx_1^2} = \frac{d}{dx_1} \left(-\frac{u_1}{u_2} \right) = \frac{-u_2 \frac{du_1}{dx_1} + u_1 \frac{du_2}{dx_1}}{u_2^2} \dots\dots\dots(1.3)$$

And

$$u_1 = u_1(x_1, x_2)$$

$$du_1 = \frac{\partial u_1}{\partial x_1} dx_1 + \frac{\partial u_1}{\partial x_2} dx_2$$

$$du_1 = u_{11} dx_1 + u_{12} dx_2$$

Therefore

$$\frac{du_1}{dx_1} = u_{11} + u_{12} \frac{dx_2}{dx_1}$$

Applying same result to u_2 , one obtains that

$$u_2 = u_2(x_1, x_2)$$

$$du_2 = \frac{\partial u_2}{\partial x_1} dx_1 + \frac{\partial u_2}{\partial x_2} dx_2$$

$$du_2 = u_{12}dx_1 + u_{22}dx_2$$

$$\frac{du_2}{dx_1} = u_{12} + u_{22} \frac{dx_2}{dx_1}$$

Substituting these results, one obtains the following:

$$\begin{aligned} \frac{d^2x_2}{dx_1^2} &= \frac{-u_2 \frac{du_1}{dx_1} + u_1 \frac{du_2}{dx_1}}{u_2^2} = \frac{-u_2 \left(u_{11} + u_{12} \frac{dx_2}{dx_1} \right) + u_1 \left(u_{12} + u_{22} \frac{dx_2}{dx_1} \right)}{u_2^2} \\ &= \frac{-u_2 u_{11} - u_2 u_{12} \left(-\frac{u_1}{u_2} \right) + u_1 u_{12} + u_1 u_{22} \left(-\frac{u_1}{u_2} \right)}{u_2^2} \\ &= \frac{-u_2 u_{11} + \frac{u_2 u_{12} u_1}{u_2} + u_1 u_{12} - \frac{u_1^2 u_{22}}{u_2}}{u_2^2} \\ &= -\frac{1}{u_2^3} \left(u_2^2 u_{11} - u_1 u_2 u_{12} - u_1 u_2 u_{12} + u_1^2 u_{22} \right) \\ &= -\frac{1}{u_2^3} \left(u_2^2 u_{11} - 2u_1 u_2 u_{12} + u_1^2 u_{22} \right) \end{aligned}$$

On substituting

$$u_1 = \frac{p_1 u_2}{p_2}$$

$$\begin{aligned} \frac{d^2x_2}{dx_1^2} &= -\frac{1}{u_2^3} \left(u_1 u_2^2 - 2u_{12} \frac{p_1 u_2}{p_2} u_2 + u_{22} \frac{p_1^2 u_2^2}{p_2^2} \right) \\ &= -\frac{u_2^2}{u_2^3 p_2^2} \left(u_{11} p_2^2 - 2u_{12} p_1 p_2 + u_{22} p_1^2 \right) \\ &= -\frac{1}{u_2 p_2^2} \left(u_{11} p_2^2 + u_{22} p_1^2 - 2u_{12} p_1 p_2 \right) \end{aligned}$$

and since the expression in the brackets is less than zero (proved above), then

$$\frac{d^2x_2}{dx_1^2} > 0$$

The above result can be proved for any number of variables.

1.7 Elasticities

Suppose one has derived a demand equation for a country, and he wants to compare this to another demand equation obtained from a different country. Sometimes, comparisons of the demand functions across commodities or even across countries, where prices vary (or there is different currency) is needed. The economists decided to introduce special units called elasticities, which are independent of any unit of measurement. These can be defined directly from the demand function. The elasticities of demand are convenient summaries of the responsiveness of the quantity demanded to factors influencing it, partly because they are independent of the units of measurement of the good, prices, income and so on. The most used elasticities and their function are listed below.

1. The *own price elasticity of demand* for good i is $\varepsilon_i = \frac{\partial \ln x_i}{\partial \ln p_i}$. This gives the percentage change in the quantity demanded for a 1% change in the price of this good. The good is said to be *price elastic* if $|\varepsilon_i| > 1$ and *price inelastic* if $|\varepsilon_i| < 1$.
2. The *income elasticity of demand* for good i is $\eta_i = \frac{\partial \ln x_i}{\partial \ln y}$. This gives the percentage change in the quantity demand for a 1% change in income. This is negative for an inferior good, and positive for a superior good. The good is said to be *income elastic* if $\eta_i > 1$ and *income inelastic* if $\eta_i < 1$.
3. The *cross-price elasticities* of demand are $\varepsilon_{ij} = \frac{\partial \ln x_i}{\partial \ln p_j}$. This gives the effect of a change in the price of one good on the demand for the other good. Obviously, one can see that $\varepsilon_{ii} = \varepsilon_i$ and that $\varepsilon_{ij} \neq \varepsilon_{ji}$.

1.8 Constraints for Demand Functions

Economic theory suggests that the demand equations obtained from maximising the utility function must satisfy certain restrictions known as constraints. These constraints will be used in chapter 2 to show that the models that the authors created for demand equations satisfy these constraints.

The constraints that a demand equation must satisfy are given below:

1. Budget Constraint

The demand equations must satisfy the budget constraint i.e.

$$\sum_i p_i x_i = y$$

The constraint can also be expressed as

$$\sum_i s_i = 1$$

where $s_i = \frac{p_i x_i}{y}$ are called the budget shares defined as the proportion of income spent on each of the goods.

2. Homogeneity

Every demand equation must be homogenous of degree zero in income and prices; i.e. if all prices and income are multiplied by a positive constant k , the quantity demanded must remain unchanged. Using Euler's theorem of homogeneity, this can be written as

$$\sum_j \frac{\partial x_i}{\partial p_j} p_j + \frac{\partial x_i}{\partial y} y = 0$$

This is expressed in elasticity form as

$$\sum_j \varepsilon_{ij} + \eta_i = 0$$

3. The Slutsky Conditions

The negative semidefiniteness conditions, i.e.

$$\frac{\partial x_i}{\partial p_i} + \frac{\partial x_i}{\partial y} x_i \leq 0$$

written in elasticity form as

$$\varepsilon_{ii} + s_i \eta_i \leq 0.$$

The symmetry conditions, i.e.

$$\frac{\partial x_i}{\partial p_j} + \frac{\partial x_j}{\partial p_i} x_j = \frac{\partial x_j}{\partial p_i} + \frac{\partial x_i}{\partial p_j} x_i$$

written in elasticity form as

$$\sum_j \frac{\varepsilon_{ij}}{s_j} + \eta_i = \sum_i \frac{\varepsilon_{ji}}{s_i} + \eta_j$$

4. The Aggregation Conditions

The Engel aggregation condition, i.e.

$$\sum_i p_i \frac{\partial x_i}{\partial y} = 1$$

This states that the sum of the marginal expenditures is unity. It is also known as the *adding-up* condition. In terms of elasticities, it can be written as

$$\sum_i s_i \eta_i = 1$$

The Cournot aggregation conditions, i.e.

$$\sum_j p_j \frac{\partial x_j}{\partial p_i} + x_i = 0$$

In terms of elasticities

$$\sum_j \frac{s_j}{s_i} \varepsilon_{ji} + 1 = 0$$

But how can we go back to preferences from a give set of demand equations? The conditions that will allow us to do so are known as the integrability problem. However, once the maximisation of utility is treated as a minimisation of costs, it becomes apparent that the demand functions must allow integration into a concave, linearly homogeneous function, that is the cost function. Thus the conditions needed for demand functions to be consistent with preferences are known as integrability

conditions. In practice, the demand functions that we are presented with are Marshallian demand functions, with total expenditure and prices as arguments, while it is Hicksian demand functions that are the derivatives of the cost functions and that integrate into it.

For suppose that we have a system of Marshallian demand functions

$$x_i = g_i(y, p)$$

By using property 4 in section 1.5, this is transformed into

$$\frac{\partial c}{\partial p_i} = g_i(c, p)$$

which is a system of partial differential equations to be solved for c as a function of p . Since integration is carried out along an indifference surface, utility is unchanged during integration and will emerge as a constant of integration. However, one should note that not all such systems of partial differential equations have a solution. This is because not all demand functions satisfy adding-up, homogeneity, symmetry and negativity.

The conditions that ensure that a function that solves $\frac{\partial c}{\partial p_i} = g_i(c, p)$ exists is

$$\frac{\partial g_i}{\partial q} g_j(q, p) + \frac{\partial g_j}{\partial p_j} g_i(q, p) = \frac{\partial g_j}{\partial q} g_i(q, p) + \frac{\partial g_i}{\partial p_i} g_j(q, p)$$

But the left-hand side of the formula is simply the formula for calculating the substitution term s_{ij} . Hence the above formula tells us that the symmetry of the substitution matrix is the fundamental integrability condition of demand theory. For the function that results from integration to be a proper cost function, it must be concave and linearly homogeneous. This will happen automatically if the substitution matrix is symmetric, negative semi-definite and satisfies

$$\sum_k s_{ik} p_k = 0$$

Then, one can say that utility maximisation results in demand functions that add up, are homogeneous of degree 0 and have symmetric negative semi-definite compensated price responses.

Conversely, demand functions that add up are homogeneous of degree 0 and have symmetric negative semi-definite compensated price responses are integrable into a consistent preference ordering. This tells us that the conditions derived in the previous section are the only consequences of utility maximisation.

The above conditions represent the principal content of the classical theory of the consumer of two goods. Each of the general restrictions defines a relationship that is exact. As these are derived from theoretical arguments, there is no reason why measured behaviour should obey them, as theory is always a simplification of reality. One of the reasons that these restrictions will not be obeyed is that statistical data contains some errors of measurement. What one can hope is that rough estimate, computed without imposing the constraints, will be not inconsistent with them.

Chapter 2

MODELLING HOUSEHOLD BEHAVIOUR

Chapter 2

MODELLING HOUSEHOLD BEHAVIOUR

2.1 *Models for Demand Equations*

The theory presented in chapter 1 offers a framework within which one can attempt to organise and interpret the data. This chapter discusses the application of some models, each of which seeks to describe the way in which income is spent. These models not only show how the theory can be made specific, but also allow us to check the usefulness of the theory in empirical work. The models will be presented in rough chronological order. The first part of the chapter is devoted to the ‘Linear Expenditure System’ (LES), the first practical model to be based entirely upon the theory. Since this model is the heart of all the models, the full theory of the LES shall be reproduced in this chapter.

Then a review of other models derived from the LES, namely the ‘Rotterdam Model’, the ‘Transcendental Logarithmic Model’, the ‘Almost Ideal Demand System’, and ‘An implicitly additive Demand System’ is carried out. For the last four models, the authors estimated a general functional form with and without the restrictions, thus testing whether the restrictions were validly imposed. In that section, I shall first state the functional form used by the authors, then review any results they obtained from their testing.

2.1.1 The Linear Expenditure System

Consider the Stone-Geary utility function³⁷

$$u = \sum_i \beta_i \log(x_i - \gamma_i),$$

where β_i is the propensity to consume
 x_i is the expenditure of the i^{th} commodity
 γ_i is a constant

subject to the linear constraint

$$\sum_i p_i x_i = y$$

where p_i is the price of the i^{th} commodity
 y is the total consumer expenditures

For normalising purposes, the condition that $\sum_i \beta_i = 1$ is also assumed.

The Lagrangian function for the above constraint is

$$L = u - \lambda(\sum_i p_i x_i - y).$$

Thus on differentiating, the following differential equations are obtained:

$$\frac{\partial L}{\partial x_i} = \frac{\partial u}{\partial x_i} - \lambda p_i$$

$$\frac{\partial L}{\partial \lambda} = -\sum_i p_i x_i + y$$

For a maximum, $\frac{\partial L}{\partial x_i}, \frac{\partial L}{\partial \lambda}$ are equal to zero.

Hence the following equations are obtained:

$$\frac{\partial u}{\partial x_i} = \lambda p_i$$

$$y = \sum_i p_i x_i$$

But the utility function was given to be:

$$u = \sum_i \beta_i \log(x_i - \gamma_i)$$

Hence

$$\frac{\partial u}{\partial x_i} = \frac{\beta_i}{x_i - \gamma_i} = \lambda p_i$$

$$\beta_i = \lambda p_i (x_i - \gamma_i)$$

But

$$\sum_i \beta_i = 1$$

$$1 = \lambda \sum_i \{p_i (x_i - \gamma_i)\}$$

$$\lambda = \frac{1}{\sum_i \{p_i (x_i - \gamma_i)\}}$$

Using the equation $y = \sum_i p_i x_i$, the following is obtained.

$$\lambda = \frac{1}{y - \sum_i p_i \gamma_i}$$

$$\beta_i = \frac{p_i (x_i - \gamma_i)}{y - \sum_j p_j \gamma_j}$$

$$\beta_i y - \beta_i \sum_j p_j \gamma_j = p_i x_i - p_i \gamma_i$$

$$p_i x_i = p_i \gamma_i + \beta_i y - \beta_i \sum_j p_j \gamma_j$$

$$x_i = \gamma_i + \frac{\beta_i}{p_i} \left[y - \sum_j p_j \gamma_j \right]$$

.....(2.1)

The last result gives the Linear Expenditure System.

2.1.1.1 Testing the restrictions

To show that the Linear Expenditure System is in fact a valid demand equation, one has to prove the four constraints of the demand functions already mentioned in the last section of chapter 1.

1. *Budget Constraint.*

This axiom is assumed true throughout the model.

2. *Homogeneity.*

Equation (2.1) is re-written in the following format:

$$x_i = \gamma_i + \frac{\beta_i y}{p_i} - \beta_i \sum_j \frac{p_j \gamma_j}{p_i}$$

and simplified to

$$x_i = \sum_j \alpha_{ij} \frac{p_j}{p_i} + \beta_i \frac{y}{p_i} \dots\dots\dots(2.2)$$

where

$$\sum_j \alpha_{ij} \frac{p_j}{p_i} = \gamma_i - \beta_i \sum_j \frac{p_j \gamma_j}{p_i}$$

But to prove homogeneity, the following must be proved

$$\sum_j p_j \frac{\partial x_i}{\partial p_j} + y \frac{\partial x_i}{\partial y} = 0 \dots\dots\dots(2.3)$$

Using equation (2.2) and partial differentiating, the following are obtained

$$\frac{\partial x_i}{\partial p_i} = -\frac{1}{2} \sum_{i \neq j} \alpha_{ij} p_j - \frac{\beta_i y}{p_i^2}$$

$$\frac{\partial x_i}{\partial p_j} = \frac{\alpha_{ij}}{p_i}$$

$$\frac{\partial x_i}{\partial y} = \frac{\beta_i}{p_i}$$

On substituting these results in (2.3), the following is obtained

$$\sum_j \frac{p_i}{x_i} \left[-\frac{1}{2} \sum_{i \neq j} \alpha_{ij} p_j \right] - \sum_{j \neq i} \frac{p_i}{x_i} \frac{\beta_i y}{p_i^2} + \sum_{j \neq i} \frac{p_i}{x_i} \frac{\alpha_{ij}}{p_i}$$

$$-\frac{1}{p_i x_i} \sum_{j \neq i} \alpha_{ij} p_j - \frac{\beta_i y}{x_i p_i} + \frac{1}{p_i x_i} \sum_{j \neq i} p_j \alpha_{ij} = -\frac{\beta_i y}{x_i p_i}$$

But

$$-\frac{y}{x_i} \frac{\partial x_i}{\partial y} = -\frac{y}{x_i} \frac{\beta_i}{p_i}$$

(QED)

3. The Slutsky Conditions

The negative semi-definiteness condition,

$$\frac{\partial x_i}{\partial p_i} + \frac{\partial x_i}{\partial y} x_i \leq 0.$$

Since

$$x_i = \gamma_i + \frac{\beta_i}{p_i} \left[y - \sum_j p_j \gamma_j \right]$$

from 2.1

$$\frac{\partial x_i}{\partial p_i} = -\frac{\beta_i y}{p_i^2} + \frac{\beta_i}{p_i^2} \sum_{j \neq i} p_j \gamma_j$$

then

$$\begin{aligned} \frac{\partial x_i}{\partial p_i} + \frac{\partial x_i}{\partial y} x_i &= -\frac{\beta_i y}{p_i^2} + \frac{\beta_i}{p_i^2} \sum_{j \neq i} p_j \gamma_j + \frac{\beta_i x_i}{p_i} \\ &= -\frac{\beta_i y}{p_i^2} + \frac{\beta_i}{p_i^2} \left[\sum_j p_j \gamma_j - p_i \gamma_i \right] + \frac{\beta_i x_i}{p_i} \\ &= -\frac{\beta_i}{p_i^2} \left[y - \sum_j p_j \gamma_j \right] - \frac{\beta_i \gamma_i}{p_i} + \frac{\beta_i x_i}{p_i} \end{aligned}$$

Arranging equation 2.1, one gets

$$\frac{x_i - \gamma_i}{p_i} = \frac{\beta_i}{p_i^2} \left[y - \sum_j p_j \gamma_j \right]$$

$$\begin{aligned}
& -\frac{\beta_i}{p_i^2} \left[y - \sum_j p_j \gamma_j \right] - \frac{\beta_i \gamma_i}{p_i} + \frac{\beta_i x_i}{p_i} \\
& - \frac{x_i - \gamma_i}{p_i} + \frac{\beta_i}{p_i} (x_i - \gamma_i) \\
& \left(\frac{x_i - \gamma_i}{p_i} \right) (\beta_i - 1)
\end{aligned}$$

For the negative condition to be satisfied, the last expression obtained must be less than zero. But $p_i > 0$ and $x_i > \gamma_i$ for all values. For the above expression to be satisfied, thus $\beta_i - 1 < 0$. This can only be true if $0 < \beta_i < 1$. So, the above restriction will be imposed on the linear expenditure system, to make sure that it satisfies the first Slutsky condition.

The symmetry conditions,

$$\frac{\partial x_i}{\partial p_j} + x_j \frac{\partial x_i}{\partial y} = \frac{\partial x_j}{\partial p_i} + x_i \frac{\partial x_j}{\partial y} \quad \dots\dots\dots(2.4)$$

where $i \neq j$

But from (2.2)

$$x_i = \sum_j \alpha_{ij} \frac{p_j}{p_i} + \beta_i \frac{y}{p_i} .$$

Therefore:

$$\begin{array}{ccc}
\frac{\partial x_i}{\partial p_j} = \frac{\alpha_{ij}}{p_i} & & \frac{\partial x_j}{\partial p_i} = \frac{\alpha_{ji}}{p_j} \\
\frac{\partial x_i}{\partial y} = \frac{\beta_i}{p_i} & \text{and} & \frac{\partial x_j}{\partial y} = \frac{\beta_j}{p_j}
\end{array}$$

Hence, on substituting:

$$\frac{\alpha_{ij}}{p_i} + x_j \frac{\beta_i}{p_i} = \frac{\alpha_{ji}}{p_j} + x_i \frac{\beta_j}{p_j}$$

where $i \neq j$

$$\alpha_{ij} p_j + x_j \beta_i p_j = \alpha_{ji} p_i + x_i \beta_j p_i$$

$$\alpha_{ij} p_j + \beta_i \left[\sum_{i=1}^n \alpha_{ji} p_i + \beta_j y \right] = \alpha_{ji} p_i + \beta_j \left[\sum_{i=1}^n \alpha_{ji} p_i + \beta_i y \right]$$

$$\alpha_{ij} p_j + \beta_i \sum_{i=1}^n \alpha_{ji} p_i + \beta_i \beta_j y = \alpha_{ji} p_i + \beta_j \sum_{i=1}^n \alpha_{ji} p_i + \beta_j \beta_i y$$

$$\alpha_{ij} p_j + \beta_i \left[\alpha_{ji} p_i + \alpha_{jj} p_j + \sum_{i=1}^n \alpha_{ji} p_i \right] = \alpha_{ji} p_i + \beta_j \left[\alpha_{ii} p_i + \alpha_{ij} p_j + \sum_{i=1}^n \alpha_{ii} p_i \right]$$

Combining together similar terms, the following is obtained:

$$\left[\alpha_{ij} + \beta_i \alpha_{jj} \right] p_j + \beta_i \alpha_{ji} p_i + \beta_i \sum_{k=1}^n \alpha_{jk} p_k = \left[\alpha_{ji} + \beta_j \alpha_{ii} \right] p_i + \beta_j \alpha_{ij} p_j + \beta_j \sum_{k=1}^n \alpha_{ik} p_k$$

where $k \neq i \neq j$

Thus, on comparing coefficients,

$$\alpha_{ij} + \beta_i \alpha_{jj} = \beta_j \alpha_{ij}$$

$$\alpha_{ij} = \beta_i \frac{(\alpha_{jj})}{(\beta_j - 1)}$$

$$\alpha_{ji} \beta_i = \alpha_{ji} + \beta_j \alpha_{ii}$$

$$\alpha_{ii} = \frac{\alpha_{ji} \beta_i - \alpha_{ji}}{\beta_j}$$

$$\alpha_{ii} = \beta_i \left[\frac{\alpha_{ji}}{\beta_j} \right] - \left[\frac{\alpha_{ji}}{\beta_j} \right]$$

and

$$\beta_i \sum_{k=1}^n \alpha_{jk} = \beta_j \sum_{k=1}^n \alpha_{ik}$$

$$\Rightarrow \alpha_{ik} = \beta_i \frac{\alpha_{jk}}{\beta_j}$$

where $k \neq i \neq j$

Defining

$$\begin{aligned} \gamma_k &= \frac{\alpha_{jk}}{\beta_j} & k \neq j \\ \gamma_k &= \frac{\alpha_{jk}}{(\beta_j - 1)} & k = j \end{aligned} \quad \text{for}$$

the following general formula is obtained :

$$\begin{aligned} \alpha_{ik} &= \beta_i \gamma_k - \delta_{ik} \gamma_i \\ i &= 1..n \\ k &= 1..n \end{aligned}$$

where

$$\begin{aligned} \delta_{ik} &= 0 & i \neq k \\ \delta_{ik} &= 1 & i = k \end{aligned} \quad \text{for}$$

So, to have a demand equation that satisfies the symmetry conditions, one has to define the price coefficients of $(n-1)$ equations in terms of the price and income coefficients of the j^{th} equation and its own income coefficient.

4. The Aggregation Conditions

The Engel aggregation condition, i.e.

$$\sum_i p_i \frac{\partial x_i}{\partial y} = 1$$

But, from above

$$\frac{\partial x_i}{\partial y} = \frac{\beta_i}{p_i}$$

So,

$$\sum_i p_i \frac{\partial x_i}{\partial y} = \sum_i p_i \frac{\beta_i}{p_i} = \sum_i \beta_i = 1$$

by assumption.

The Cournot aggregation Condition, i.e.

$$\sum_j p_j \frac{\partial x_j}{\partial p_i} + x_i = 0.$$

But since

$$x_i = \gamma_i + \frac{\beta_i}{p_i} y - \frac{\beta_i}{p_i} \sum_j p_j \gamma_j$$

$$\frac{\partial x_i}{\partial p_i} = -\frac{\beta_i y}{p_i^2} + \frac{\beta_i}{p_i^2} \sum_{j \neq i} p_j \gamma_j$$

and

$$\frac{\partial x_i}{\partial p_j} = -\frac{\beta_i}{p_i} \gamma_j$$

then

$$\begin{aligned} & \sum_j p_j \frac{\partial x_j}{\partial p_i} + x_i \\ & p_i \frac{\partial x_i}{\partial p_i} + \sum_{j \neq i} p_j \frac{\partial x_j}{\partial p_i} + x_i \\ & p_i \left[-\frac{\beta_i y}{p_i^2} + \frac{\beta_i}{p_i^2} \sum_{j \neq i} p_j \gamma_j \right] + \sum_{j \neq i} p_j \left[-\frac{\beta_j}{p_j} \gamma_i \right] + x_i \\ & -\frac{\beta_i y}{p_i} + \frac{\beta_i}{p_i} \sum_{j \neq i} p_j \gamma_j + \sum_{j \neq i} -\beta_j \gamma_i + \gamma_i + \frac{\beta_i y}{p_i} - \frac{\beta_i}{p_i} \sum_j p_j \gamma_j \\ & \frac{\beta_i}{p_i} \sum_{j \neq i} p_j \gamma_j + \sum_{j \neq i} -\beta_j \gamma_i + \gamma_i - \frac{\beta_i}{p_i} \left[p_i \gamma_i + \sum_{j \neq i} p_j \gamma_j \right] \\ & \frac{\beta_i}{p_i} \sum_{j \neq i} p_j \gamma_j - \gamma_i \sum_{j \neq i} \beta_j + \gamma_i - \beta_i \gamma_i + \frac{\beta_i}{p_i} \sum_{j \neq i} p_j \gamma_j \\ & \gamma_i - \left[\sum_j \beta_j \gamma_j \right] = 0 \end{aligned}$$

since $\sum_j \beta_j = 1$.

So, the Linear Expenditure System is in fact a valid demand equation.

However, this model also satisfies the following restrictions.

1. Given that $\beta_i > 0$, all income elasticities are positive, confirming the absence of inferior goods.

Proof :

$$\eta_i = \frac{\partial \ln x_i}{\partial \ln y}$$

$$p_i x_i = p_i \gamma_i \left[y - \sum_j p_j \gamma_j \right]$$

$$\frac{\partial \ln x_i}{\partial \ln y} = \frac{\partial x_i}{\partial y} \cdot \frac{y}{x_i} = \frac{\beta_i}{p_i} \cdot \frac{y}{x_i}$$

But p_i, y and x_i are positive, so η_i (income elasticity) is positive if $\beta_i > 0$.

2. Given that $\gamma_i > 0$, the own price elasticity of good i is smaller than one, so that good i is price inelastic.

Proof :

$$x_i = \gamma_i + \frac{\beta_i}{p_i} y - \frac{\beta_i}{p_i} \sum_j p_j \gamma_j$$

$$\frac{\partial \ln x_i}{\partial \ln p_i} = \frac{\partial x_i}{\partial p_i} \cdot \frac{p_i}{x_i}$$

$$\left[\frac{-\beta_i}{p_i^2} y + \frac{\beta_i}{p_i^2} \sum_{j \neq i} p_j \gamma_j \right] \cdot \frac{p_i}{x_i}$$

$$\frac{-\beta_i y}{p_i x_i} + \frac{\beta_i}{p_i x_i} \sum_{j \neq i} p_j \gamma_j$$

$$\frac{-\beta_i}{p_i x_i} \left[y - \sum_{j \neq i} p_j \gamma_j \right] \dots\dots\dots(2.5)$$

But

$$p_i x_i = p_i \gamma_i + \beta_i \left[y - \sum_{j \neq i} p_j \gamma_j - p_i \gamma_i \right]$$

So (2.5) becomes

$$\begin{aligned} & \frac{-\beta_i}{p_i x_i} \left[\frac{p_i x_i - p_i \gamma_i + \beta_i p_i \gamma_i}{\beta_i} \right] \\ & \frac{-1}{x_i} [x_i - \gamma_i + \beta_i \gamma_i] \\ & -1 + \frac{\gamma_i - \beta_i \gamma_i}{x_i} \\ & -1 + \frac{\gamma_i [1 - \beta_i]}{x_i} \end{aligned}$$

which means that when $\gamma_i > 0$, then the own price elasticity is less than one.

3. When good j is price elastic, the cross-price elasticity is positive.

Proof :

$$\varepsilon_{ij} = \frac{\partial \ln x_i}{\partial \ln p_j}$$

But

$$\begin{aligned} \frac{\partial \ln x_i}{\partial \ln p_j} &= \frac{\partial x_i}{\partial p_j} \cdot \frac{p_j}{x_i} \\ &= \frac{-\beta_j \gamma_j}{p_i} \cdot \frac{p_j}{x_i} \end{aligned}$$

So when good j is price elastic, from previous proof, it follows that $\gamma_j < 0$, and since β_j, p_j, p_i and x_i are all positive, this implies that the cross-price elasticity is negative.

The Linear Expenditure System is widely used because it can be interpreted as stating that expenditure on good j , given as $p_j x_j$ can be decomposed into two components—the first is the expenditure on a certain base amount, which is the minimum expenditure to which the consumer is committed. The second part is the supernumerary income. This model is also one of the few that satisfies all nine theoretical restrictions, and it can be derived from a specific utility function, as shown above.

One of the striking features of the LES is that it has only $2n$ parameters, $(2n - 1)$ of which may be chosen independently. This means that the LES in fact is more specialised than might be needed in practice. Also it can be derived from a specific utility function. The LES is one of the few systems that satisfy all the restrictions. In fact, as one can see from the above discussion, LES satisfies all the integrability conditions. However since the model was constrained to satisfy these conditions, then these cannot be tested on the model. Besides this, it was shown that LES satisfies some other restrictions, such as the absence of inferior goods, etc. These restrictions are not desirable in the sense that there is no valid reason from the point of view of utility maximisation why these should be imposed. The adding-up problems of the linear expenditure system have always been understood but were usually thought to be unimportant since all studies considered only a fraction of the total budget. Besides, the LES belongs to a class of demand models, all of which share the property of approximate proportionality between price and expenditure elasticities. Unless this is what is needed, which in most cases, this will not be the case, the LES will be considered as too restrictive, and other models will need to be found.

In formulating the Linear Expenditure System, Stone³⁷ started from a utility function, and tried to restrict the function such that it satisfies the restrictions needed to be a demand function. An alternative way is to define a demand function and try to test

whether the restrictions hold on this function using the data obtained, that is to check whether it is a valid demand equation. Theil and Barten⁹ tried the latter approach and the model thus created is known as the Rotterdam model.

2.1.2 *The Rotterdam Model*

The starting point of the Rotterdam model is the equation $\log x_i = \alpha_i + e_i \log y + \sum_k e_{ik} \log p_k + u_i$. This is also the starting point for the LES.

However, the approach is different. In his work, Stone³⁷ imposed the restrictions on this demand equation, finally transforming it into the LES. Theil and Barten⁹ differentiate the above equation to obtain

$$d \log x_i = e_i d \log y + \sum_j e_{ij} d \log p_j$$

where the total expenditure elasticity, and the cross-price elasticity are not necessarily constant, as was assumed by Stone³⁷ in his formulation of the LES. The Slutsky decomposition is used to write $e_{ij} = e_{ij}^* - e_i w_j$, for compensated price elasticities e_{ij}^* , such that the equation becomes

$$\begin{aligned} d \log x_i &= e_i d \log y + \sum_j (e_{ij}^* - e_i w_j) d \log p_j \\ \Rightarrow d \log x_i &= e_i d \log y + \sum_j (e_{ij}^*) d \log p_j - \sum_j (e_i w_j) d \log p_j \\ \Rightarrow d \log x_i &= e_i \left(d \log y - \sum_k (w_k) d \log p_k \right) + \sum_j (e_{ij}^*) d \log p_j \end{aligned}$$

However, this equation does not readily lend itself to symmetry, so the budget shares w_i must also be included in the equation. Thus the demand equation is

$$w_i d \log x_i = b_i d \log \bar{y} + \sum_j (c_{ij}) d \log p_j$$

where

$$d \log \bar{y} = d \log y - \sum_k (w_k) d \log p_k$$

$$b_i = w_i e_i = p_i \frac{\partial x_i}{\partial y}$$

follows from the budget constraint.

It shows that b_i is the marginal propensity to consume.

$$c_{ij} = w_i e_{ij}^* = \frac{p_i p_j s_{ij}}{y}$$

where s_{ij} is the $(i,j)^{\text{th}}$ term of the Slutsky substitution matrix.

2.1.2.1 Testing the restrictions

Since no restrictions were placed on the Rotterdam model, the authors had to test whether the model defined violates the restrictions.

1. Budget Constraint

This constraint is assumed true throughout the model.

2. Homogeneity

For the Rotterdam model to be homogeneous, the following restriction has to be made to the model

$$\sum_k c_{jk} = 0$$

This can be imposed and tested equation by equation.

3. The Slutsky Conditions

The negative semi-definiteness condition

Since all the prices are positive, then to test the negative semi-definite of matrix S one needs only to test that C is negative semi-definite.

The symmetry conditions

The symmetry of the substitution matrix implies that the matrix C be symmetric.

Hence one should only show that $c_{ij} = c_{ji}$, for all i and j .

Therefore, in practice, tests for symmetry and negative semi-definite of S were tested on C .

4. *The Aggregation Conditions*

The adding up restriction requires that the marginal propensities to spend on each good sum up to unity and that the net effect of a price change on the budget be zero. In terms of the Rotterdam model, this is defined as

$$\sum_k b_k = 1 \text{ and } \sum_k c_{kj} = 0$$

In practice, the discrete approximations to the differentials are usually chosen in such a way as to ensure that the above equations are valid.

Barten⁹ was the first one to test the homogeneity and symmetry of the Rotterdam model in 1967 using Dutch data. He conducted two experiments. In the first one, he considered only four broad groups and used informal testing procedures. During that particular study, he found that his model fitted the data and satisfied the constraints for a demand equation.

In the second experiment, he used 16 groups and an explicit maximum likelihood approach. First he estimated the equation without any restriction except that of adding-up, which is embodied in the data as with all models. Then he applied the homogeneity conditions. On the effect that this restriction had on the data, Barten⁹ concluded that on his data, a proportional change in prices and aggregate expenditure

would not leave the pattern of demand unchanged. On imposing also the symmetry conditions, he again found a conflict between theory and evidence, however, three years later, Deaton¹¹ concluded that given that the homogeneity conditions had already imposed on the data, the symmetry conditions is not particularly damaging.

Deaton¹¹ performed the same test on another set of data, for a nine-group sample using maximum likelihood estimation, and he concluded the same result as Barten⁹ did.

Although the Rotterdam model invalidates the relatively unsophisticated methodology of Stone's study (let alone more sophisticated techniques), it was the first model that had a substitution matrix estimated subject only to symmetry. Thus it was possible to identify the substitutes and complements directly from the estimation, and to see whether negativity fitted the evidence.

Because of the rejections of homogeneity when dealing with the Rotterdam model, a natural question to ask is whether these rejections are specific to the model being based perhaps on an unhappy approximation or choice of functional form, or whether they are indeed a genuine reflection of reality.

However, Byron⁹ performed similar tests on the double logarithmic model using Barten's data, and here too the results were negative. This seems to suggest that the model is in fact valid, what is invalid is the data. However this is inconclusive. Loglinear models cannot satisfy the theory exactly, so the rejections of homogeneity or symmetry may simply reflect the quality of the approximation rather than any inherent property of the data.

Although the Rotterdam model is a very useful approximation, it is worth taking a moment to mention some of its drawbacks. The Rotterdam model was originally expressed in terms of differentials and it can be shown that if these satisfy Young's

theorem of equality of cross partial derivatives, then the model collapses to one of constant budget shares, which is completely at variance with evidence.

2.2 *Flexible Functional Forms*

After the Rotterdam model, a number of researchers have approved the testing problem using what have become known as “flexible functional forms”. The basic method is to approximate the indirect utility function by some specific functional form that has enough parameters to be regarded as a reasonable approximation to whatever the true unknown function may be.

Therefore the problem is reduced into maximising $\ln u$ subject to the constraint

$\sum_i p_i x_i = y$. The Lagrangian function is given by

$$L = \ln u - \lambda \left(\sum_i p_i x_i - y \right)$$

So, the following differentials are obtained

$$\frac{\partial \ln u}{\partial x_i} = \lambda p_i \quad \text{and} \quad y = \sum_i p_i x_i$$

But

$$\frac{\partial \ln u}{\partial x_i} = \frac{\partial \ln u}{\partial \ln x_i} \frac{\partial \ln x_i}{\partial x_i} = \frac{1}{x_i} \frac{\partial \ln u}{\partial \ln x_i}$$

So

$$\frac{\partial \ln u}{\partial \ln x_j} = \frac{\partial \ln u}{\partial u} \frac{\partial u}{\partial \ln x_j} = \frac{1}{u} \frac{\partial u}{\partial \ln x_j}$$

and

$$\frac{\partial u}{\partial \ln x_j} = u \frac{\partial \ln u}{\partial \ln x_j}$$

and

$$\frac{\partial u}{\partial \ln x_j} = \frac{\partial u}{\partial x} \frac{\partial x}{\partial \ln x_j} = \mu p_j \frac{\partial x}{\partial \ln x_j}$$

But

$$\frac{\partial \ln x_j}{\partial x} = \frac{1}{x_j} \Rightarrow \frac{\partial u}{\partial \ln x_j} = \mu p_j x_j$$

Hence

$$\frac{\partial \ln u}{\partial \ln x_j} = \frac{1}{u} \frac{\partial u}{\partial \ln x_j} = \frac{\mu p_j x_j}{u}$$

where μ is the marginal propensity to consume. Arranging the equation and summing across the commodities, the following equation is obtained:

$$\frac{\mu}{u} = \frac{1}{y} \sum \frac{\partial \ln u}{\partial \ln x_i}$$

such that

$$\frac{\partial \ln u}{\partial \ln x_j} = \frac{p_j x_j}{y} \sum \frac{\partial \ln u}{\partial \ln x_i} \dots\dots\dots(2.6)$$

2.2.1 *The Transcendental Logarithmic Model*

In their paper, Christenson, Jorgenson and Lau⁷ present the reader with two functions, a direct translog utility function, and an indirect translog utility function. The two functions are duals to each other. The starting utility function for the translog function is a function quadratic in the logarithms of the quantities consumed. The function used for approximation is given as the negative of the logarithm of the direct utility function. This is to preserve symmetry with the treatment of the indirect utility function. The utility function used allows expenditure shares to vary with the level of total expenditure and permit a greater variety of substitution patterns among commodities. The direct translog utility function is defined as

$$-\ln u = \alpha_0 + \sum \alpha_i \ln x_i + \frac{1}{2} \sum \sum \beta_{ij} \ln x_i \ln x_j$$

Using this form in equation (2.6), one obtains

$$\alpha_j + \sum \beta_{ji} \ln x_i = \frac{p_j x_j}{y} \sum (\alpha_k + \sum \beta_{ki} \ln x_i)$$

where $(j = 1, 2, \dots, m)$

To simplify the notation, let

$$\alpha_M = \sum \alpha_k \text{ and } \beta_{Mi} = \sum \beta_{ki}$$

where $(i = 1, 2, \dots, m)$

such that the equation becomes

$$\frac{p_j x_j}{y} = \frac{\alpha_j + \sum \beta_{ji} \ln x_i}{\alpha_M + \sum \beta_{Mi} \ln x_i}$$

where $(j = 1, 2, \dots, m)$

Since

$$\frac{\sum p_i x_i}{y} = 1$$

this would imply that given the parameters of any $m-1$ equations for the budget

shares $\frac{p_j x_j}{y}$, for $(j = 1, 2, \dots, m)$, then the parameters of the m^{th} equation α_m and β_{mj}

can be determined from the definitions of α_M and β_{Mj} . Besides, since the equations

of the budget shares are homogeneous of degree zero in the parameters, a

normalisation of the parameters is required for estimation. And the normalisation

used by Christenson, Jorgenson and Lau⁷ in their paper was $\alpha_M = \sum \alpha_i = -1$.

2.2.1.1 Testing the restrictions

1. Budget Constraint

This axiom is assumed true throughout the model.

2. Homogeneity

If u is additive, i.e.

$$u(x_1, x_2) = u(x_1) + u(x_2) \quad \dots\dots\dots(2.7)$$

then Christenson, Jorgenson and Lau⁷ showed that there exists a real-valued function of one variable F such that

$$\ln u = F\left(\sum \ln u^i(x_i)\right)$$

where each of the functions u^i depends only on one of the commodities demanded. For suppose $u(x_1, x_2) = x_1 x_2$, then by (2.7),

$$x_1 x_2 = x_1 + x_2$$

On taking logs

$$\ln(x_1 x_2) = F(\ln(x_1) + \ln(x_2))$$

showing that F is the identity function.

However if $u(x_1, x_2) = x_1 + x_2$, then by (2.7)

$$x_1 + x_2 = x_1 + x_2$$

and on taking logs

$$\ln(x_1 + x_2) = F(\ln(x_1) + \ln(x_2))$$

which is not true in general.

If the direct utility function were homothetic, then

$$\ln u = F(\ln H(x_1, x_2, \dots, x_m))$$

where the function H is homogeneous of degree one. This means that if

$$u(\lambda x_1, x_2) = \lambda u(x_1, x_2), \text{ then}$$

$$\ln u(\lambda x_1, x_2) = F(\ln H(x_1, x_2)).$$

The parameters of the translog approximation to a homothetic direct utility function would be

$$-\frac{\partial \ln u}{\partial \ln x_i} = \frac{\partial F}{\partial \ln H} \frac{\partial \ln H}{\partial \ln x_i} = \alpha_i$$

for $(i = 1, 2, \dots, m)$

and

$$-\frac{\partial^2 \ln u}{\partial \ln x_i \partial \ln x_j} = \left[\frac{\partial F}{\partial \ln H} \frac{\partial^2 \ln H}{\partial \ln x_i \partial \ln x_j} + \frac{\partial^2 F}{\partial \ln H^2} \frac{\partial \ln H}{\partial \ln x_i} \frac{\partial \ln H}{\partial \ln x_j} \right] = \beta_{ij}$$

where $(i \neq j; i, j = 1, 2, \dots, m)$

Since function H is homogeneous of degree one, then

$$\sum \frac{\partial \ln H}{\partial \ln x_i} = 1$$

and

$$\sum_{j=1}^m \frac{\partial^2 \ln H}{\partial \ln x_i \partial \ln x_j} = 0$$

for $(i = 1, 2, \dots, m)$

If σ is defined as $\sigma = \frac{\partial^2 F}{\partial \ln H^2}$, then under homotheticity, the parameters of the

translog utility function satisfy the restrictions $\beta_{Mi} = \sigma \alpha_i$. A necessary and

sufficient condition for the translog utility function to be homothetic, is that it is

homogeneous such that $\sigma = 0$.

3. *The Slutsky Conditions*

Since the logarithm of the direct translog utility function is twice differentiable in the logarithms of the quantities consumed, then the Hessian of this function is symmetric. This will give rise to a set of $\frac{1}{2}m(m-1)$ symmetry restrictions relating the parameters of the cross-partial derivatives, namely that

$$\beta_{ij} = \beta_{ji}$$

where $(i \neq j; i, j = 1, 2, \dots, m)$

4. *The Equality Restrictions*

In the previous section, it was made clear that the $m-1$ equations subject to the normalisation could be estimated. Furthermore, if utility maximisation was used, then the parameters β_{Mj} appearing in each equation must be the same. This would result in a total of $m(m-2)$ restrictions referred to as the *equality restrictions*.

5. *The Aggregation Conditions*

We have already seen that if the direct utility function is additive, then

$$\ln u = F\left(\sum \ln u^i(x_i)\right)$$

Thus the parameters of the translog approximation to an additive direct utility function can be written as

$$-\frac{\partial \ln u}{\partial \ln x_i} = -F' \frac{\partial \ln u^i}{\partial \ln x_i} = \alpha_i$$

where $(i = 1, 2, \dots, m)$

and

$$-\frac{\partial^2 \ln u}{\partial \ln x_i \partial \ln x_j} = -F'' \frac{\partial \ln u^i}{\partial \ln x_i} \frac{\partial \ln u^j}{\partial \ln x_j} = \beta_{ij}$$

where $(i \neq j; i, j = 1, 2, \dots, m)$

The logarithmic derivatives are

$$F' = \frac{\partial \ln u}{\partial \ln u^i}$$

for $(i = 1, 2, \dots, m)$

and

$$F'' = \frac{\partial^2 \ln u}{\partial \ln u^i \partial \ln u^j}$$

for $(i, j = 1, 2, \dots, m)$

If θ is defined as $\theta = -\frac{F''}{(F')^2}$, then under additivity, the parameters of the

translog function satisfy the $\frac{1}{2}(m-2)(m-1)$ additivity restrictions $\beta_{ij} = \theta \alpha_i \alpha_j$

where $(i \neq j; i, j = 1, 2, \dots, m)$. Thus the direct translog utility function will be additive if and only if $\ln u$ can be written as the sum of m functions. Such a function is referred to as explicitly additive. Thus explicit additivity of the translog function will imply the additivity restrictions plus a further restriction, namely that $\theta = 0$. This restriction is called the *explicit additivity restriction*.

After an examination of American consumption Christenson, Jorgenson and Lau⁷ conclude that 'Our results.... Make possible an unambiguous rejection of the theory of demand'. However, this statement is not quite true, this is because there is no reason to suppose that utility functions are exactly 'translogarithmic' either for

individual consumers or in aggregate so that the approximation will be less accurate. The only guarantee one can get from this model is that it is accurate in the locality of some point, at particular levels of the price-income ratios. Therefore, like in the Rotterdam model, no one can know whether the hypothesis is false or whether the approximation was inaccurate. As argued by Deaton⁹, afterwards, the most impressive result is the uniformity of the results, because the Rotterdam, loglinear, and translog models all make different approximations and yet all obtain the same results, namely that the restrictions of the theory do not hold, at least on aggregate data.

2.2.2 *The Almost Ideal Demand System*

The demand functions in the previous two models were complicated and clumsy to estimate. Another model with properties similar to the translog but much simpler was needed. This model is in fact a generalisation of the translog model, but much simpler to use. The model, called the Almost Ideal Demand System or AIDS for short, was defined by Deaton and Muellbauer⁹ in 1980.

The authors defined a new class of preferences, known as the PIGLOG class, which permit exact aggregation over consumers. These classes are represented via the cost function, which defines the minimum expenditure necessary to attain a specific utility level at given prices. This is denoted by $c(u, p)$ for utility u and price vector p .

The PIGLOG class is defined by

$$\log c(u, p) = a(p) + ub(p)$$

where $a(p)$ and $b(p)$ are functions of prices and homogeneous of degree one.

These cost functions will always give rise to demand functions of the form

$$x_i = \alpha_i + \beta_i \log y$$

Using

$$x_i = \frac{\partial \log c}{\partial \log p_i}$$

then

$$\frac{\partial \log c}{\partial \log p_i} = \frac{\partial \log a(p)}{\partial \log p_i} + \frac{\partial \log b(p)}{\partial \log p_i} \frac{\partial \log u}{\partial \log p_i}$$

where

$$\alpha_i = \frac{\partial \log a(p)}{\partial \log p_i}, \beta_i = \frac{\partial \log b(p)}{\partial \log p_i}$$

Hence

$$x_i = \alpha_i + \beta_i \log y$$

The authors chose $a(p)$ as the function

$$a(p) = \alpha_0 + \sum_k \alpha_k \log p_k + \frac{1}{2} \sum_k \sum_l \gamma_{kl}^* \log p_k \log p_l$$

and $b(p)$ as

$$b(p) = \beta_0 \prod p_k^{\beta_k}$$

where α , β and γ^* are parameters.

The budget shares can be derived by using the above two equations and

$$x_i = \frac{\partial \log c}{\partial \log p_i}$$

Thus

$$x_i = \alpha_i + \sum_j \gamma_{ij} \log p_j + \beta_i \beta_0 \prod p_k^{\beta_k}$$

where $\gamma_{ij} = \frac{1}{2}(\gamma_{ij}^* + \gamma_{ji}^*) = \gamma_{ji}$.

On inverting the cost function to get the indirect utility function, the budget shares are obtained as a function of p and y , namely that

$$x_i = \alpha_i + \sum_j \gamma_{ij} \log p_j + \beta_i \log \left(\frac{x}{P} \right) \quad \dots\dots\dots(2.8)$$

where P is a price index defined by the equation

$$\log P = \alpha_0 + \sum_k \alpha_k \log p_k + \frac{1}{2} \sum_k \sum_l \gamma_{kl} \log p_k \log p_l \quad \dots\dots\dots(2.9)$$

The restrictions which the authors impose such that the system becomes a valid demand equation, are the usual constraints, i.e. the adding up restrictions,

$$\sum_{i=1}^n \alpha_i = 1, \sum_{i=1}^n \gamma_{ij} = 0, \sum_{i=1}^n \beta_i = 0$$

the homogeneity restrictions,

$$\sum_j \gamma_{ij} = 0$$

and the slusky conditions

$$\gamma_{ij} = \gamma_{ji}$$

The matrix whose elements are the γ_{ij} is not required to be negative semi-definite;

thus the negativity conditions are satisfied if the matrix C is defined by

$$c_{ij} = \gamma_{ij} + \beta_i \beta_j \log \left(\frac{q}{P} \right) - x_i \delta_{ij} + x_i x_j$$

where δ_{ij} is the Kronecker delta.

If these restrictions hold, the demand system created will add up to total expenditure, i.e. $\sum_i x_i = 1$, will be homogeneous of degree zero in prices and total expenditure taken together, and will satisfy the Slutsky symmetry conditions.

The model preserves the generality of both Rotterdam and Translog models. Also the theoretical restrictions apply directly to the parameters. The above restrictions are all implied by utility maximisation. However, unrestricted estimation of the model will only automatically satisfy the adding up restrictions such that the AIDS also offers the opportunity of testing homogeneity and symmetry by imposing the restrictions. The most interesting feature of the AIDS is that the demand equation (2.8) is very close to being linear in its parameters. Apart from expression P , the demand system can also be estimated equation by equation using ordinary least squares. As to P , the restrictions on α and γ ensure that equation (2.9) defines P as a linearly homogeneous function of the individual prices. Usually prices are relatively collinear, thus P will be approximately proportional to any appropriately defined price index, like for example $\log P = \sum x_k \log p_k$. Such an index can be calculated directly before estimation such that the demand equation will be straightforward to estimate, in contrast with the translog model. The β parameters in the AIDS determine whether the goods are luxuries or necessities, namely if $\beta_i > 0$, then the commodity is a luxury and vice versa.

2.2.2.1 *Testing the Restrictions*

The authors tested the model on British annual data, they got results better than the LES¹⁰.

1. *Budget Constraint*

This axiom is assumed true throughout the model.

2. *Homogeneity*

Deaton and Muellbauer⁹ imposed the restriction $\sum_j \gamma_{ij} = 0$ on the data.

Homogeneity was once again rejected. But they found out that for every commodity group where homogeneity was unacceptable by its F -ratio, there was a sharp drop in the Durbin-Watson statistic. This suggested that the rejection of homogeneity might be caused by an inappropriate specification of the dynamics of behaviour. For example, expenditure on several items may be relatively inflexible in the short run, housing is one of the obvious cases. The authors go on to say that no satisfactory model has been defined that explains these rejections and the aggregate consumption functions.

However, Blundell² in 1988 found that using the AIDS, the homogeneity is acceptable across all goods. He even suggests that dynamic mis-specification is the root cause of homogeneity rejections in the Deaton and Muellbauer⁹ paper.

3. *The Slutsky conditions*

The authors applied the symmetry restrictions, i.e. $\gamma_{ij} = \gamma_{ji}$ on their model.

They had no clear answer as to whether the drop in likelihood was due to the fact of the rejection of homogeneity or not.

The negativity conditions cannot be imposed globally and therefore the condition can only be tested locally.

4. *The Aggregation conditions*

Since the data add up by construction, there is no need to test this restriction.

The model gives an empirically successful way of allowing marginal budget shares to respond to income levels. But besides being defined only in differentials, the model has the problem that under large changes in real incomes, budget shares can stray outside the $[0,1]$ interval. In the AIDS, the budget shares of the various commodities are linearly related to the logarithm of real total expenditure and the logarithms of relative prices. However, since no time trends are being considered, the model rejected the homogeneity restrictions. This seems to suggest that demand functions fitted to aggregate time series data are not homogeneous and probably not symmetric.

2.2.3 *Conclusions for the above three models*

As a conclusion, three different models each having three different approximations, and fitted to three different data sets, yielded the same results, namely that demand functions fitted to aggregate time-series data are not homogeneous and probably not symmetric.

However, some problems occurring in all three models were completely ignored by all the authors. The problems are listed here.

1. The theory deals with a finite number of commodities, however all of the three models were estimated over a large group of commodities, as in fact there are thousands of commodities to consider.
2. The studies in this chapter treated aggregate data as if they had related to a single consumer. There is no general reason to suppose that this is valid.

3. It is not true that total expenditure is given exogenously and must be spent. Income is a choice variable, and so may be affected by the prices of commodities.
4. There was lack of treatment for durable goods. In some models, these are treated as a category like the others.
5. The models all have Engel curve rank equal to two, i.e. the models are not flexible enough to describe real world data adequately.

2.2.4 *An Implicitly Additive Demand System*

Rimmer and Powell³³ worked on model combining the properties of the AIDS and the LES. Besides, as proven by Lewbel²⁷ in his paper in 1991, the rank of realistic demand systems is three, while that of the previous models was two. The rank is defined as the minimum number of price indices in which cost functions can be written. Rimmer and Powell³³ later confirmed this result in 1994. The new model, called An Implicitly Additive Demand System or AIDADS for short is a rank three model, which allows the marginal budget shares to vary as a function of total expenditure.

Rimmer and Powell³³ start by defining the direct utility functions as

$$U_i = \varphi_i \ln \left(\frac{x_i - \gamma_i}{Ae^u} \right)$$

where

$$\varphi_i = \frac{[\alpha_i + \beta_i G(u)]}{[1 + G(u)]}$$

for $(i = 1, 2, \dots, n)$.

The function $G(u)$ is any positive, monotonic, twice-differentiable function.

Besides,

$$0 \leq \alpha_i, \beta_i \leq 1$$

and

$$\sum_{i=1}^n \alpha_i = 1 = \sum_{i=1}^n \beta_i$$

The authors choose $G(u) = e^u$, for the monotonicity properties of the function. Solving the utility function by using the Lagrange multipliers as shown in the LES, the following is obtained:

$$\frac{\partial u}{\partial x_i} = \lambda p_i \text{ and } y = \sum_i p_i x_i$$

On differentiating,

$$\frac{\partial u}{\partial x_i} = \frac{[\alpha_i + \beta_i G(u)]}{[1 + G(u)]} \frac{1}{(x_i - \gamma_i)} = \lambda p_i$$

Therefore on arranging the equation, the following is obtained:

$$\lambda^{-1} p_i (x_i - \gamma_i) = \frac{[\alpha_i + \beta_i G(u)]}{[1 + G(u)]} \dots\dots\dots(2.10)$$

The equation is summed across i , thus

$$\sum_{i=1}^n \lambda^{-1} p_i x_i - \sum_{i=1}^n \lambda^{-1} p_i \gamma_i = \frac{1}{[1 + G(u)]} \sum_{i=1}^n [\alpha_i + \beta_i G(u)]$$

Using

$$y = \sum_i p_i x_i, \text{ and } \sum_{i=1}^n \alpha_i = 1 = \sum_{i=1}^n \beta_i$$

the equation is transformed to

$$\lambda^{-1} y - \lambda^{-1} \sum_{i=1}^n p_i \gamma_i = \frac{1}{[1 + G(u)]} [1 + G(u)]$$

or

$$y - \sum_{i=1}^n p_i \gamma_i = \lambda$$

Denote

$$p' \gamma = \sum_{i=1}^n p_i \gamma_i$$

Thus

$$y - p' \gamma = \lambda$$

Substituting this into equation (2.10), the following is obtained:

$$\frac{p_i(x_i - \gamma_i)}{y - p' \gamma} = \frac{[\alpha_i + \beta_i G(u)]}{[1 + G(u)]} = \phi_i$$

Therefore

$$p_i(x_i - \gamma_i) = \phi_i(y - p' \gamma)$$

for $(i = 1, 2, \dots, n)$

The latter gives the AIDADS expenditure system. As a side note, remember that the LES was

$$x_i = \gamma_i + \frac{\beta_i}{p_i} \left[y - \sum_j p_j \gamma_j \right]$$

On setting every α_i equal to the β_i in ϕ_i , will reduce the equation to β_i . This will result into

$$p_i(x_i - \gamma_i) = \beta_i(y - p' \gamma)$$

that is

$$p_i x_i = p_i \gamma_i + \beta_i(y - p' \gamma)$$

which is the LES. The richer Engel possibilities in AIDADS will add $(n-1)$ parameters, that is the independent values of α_i .

2.2.4.1 *Testing the Restrictions*

Rimmer and Powell³³ used annual time series data for 35-period fiscal years 1954-55 up to 1988-89 in Australia. They tested the model on 6 commodities, namely food, alcohol and tobacco, clothing and footwear, durables, rent and other. By using the definition that if $\alpha_i > \beta_i$, then the commodity is a necessity, they found out that the first four commodities were necessities and the other luxuries. Except for the durables, this agrees with the usual classification. However, when the authors compare their values with other values obtained by other authors using different models and Australian data of the 1984, they found that the results for food commodity agrees with the results of the other authors, but there was a great difference for the other commodities. This may be due to the fact that consumption patterns do change over time, as such trends were not accounted for in this model. However, the authors still argue that their results may be better due to the fact that they are using a rank 3 model. Since the model is a direct generalisation of the LES, it was assumed that the constraints for the demand equations are satisfied.

All the above models use a micro demand equation to model the macroeconomic behaviour. This is because, up to now, no model for macro behaviour was available. However, in the next chapter we shall see that in fact, no model is needed since cross sectional data can be used.

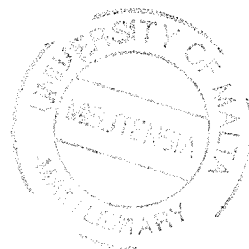
2.3 *Conclusion*

In principle, it is possible to apply the microeconomic theory of the preceding sections to the data. However, most of the time the data is available only for aggregates of households, so there are no obvious reasons why the theory, which was formulated for individual households should be directly applicable. This transition from the microeconomics to the analysis of market behaviour is known as the aggregation problem. Thus aggregation theory will provide us with rules to treat consumer behaviour as if it were the outcome of the decisions of a single maximising consumer. However, some economists have argues that the variations in circumstances of individual households average out to negligible proportions in aggregate.

In this chapter, a review of some of the existing models for demand equations has been carried out. The demand equations were obtained by maximising an objective function under some constraints. However, the choice of the objective function and the choice of the system of constraints were only subjective to the author of the model. Besides, studying optimisation will not yield any information of how human beings actually behave. Therefore, in the next chapter, an extensive analysis of how households really behave will be carried out.

Chapter 3

DENSITY ESTIMATION AND CURVE FITTING



Chapter 3

DENSITY ESTIMATION AND CURVE FITTING

Although there exists quite a number of demand functions, I have found out that no applied study was done on the Household Budgetary Survey for Malta. The HBS is a survey given to 2722 households and is carried out for all tested year. The survey considered here was done in 1994.

A problem coming from economics, that can be solved statistically is a problem associated with the data. The data obtained from surveys is considered to be non-realistic data, mainly because people tend to exaggerate their expenditure, and minimise their income. To solve the above problem, one needs to look into density estimation. Having obtained all the equations that fit the commodities, one could simulate the households such that further studies could be carried out.

In the first part of this chapter, current methods being used for density estimation both parametric and non-parametric will be reviewed, with some of these methods being tried out on Maltese data. In the second part of the chapter, a curve is fitted to each commodity and a goodness of fit test is applied. Finally a detailed study of the outliers is given after each commodity.

To study the distributions of expenditure and household commodities, parametric tests were used. This is because non-parametric analysis supplies the graphical presentation of the data without supplying the actual curve. Also, parametric methods such as least squares and maximum likelihood estimation were not suitable in this research. Although no scientific method was supplied for estimation of the data,

curves were fitted to the data points empirically, and unless otherwise proven, these appeared to be the best fit to the data points.

3.1 *Basic Definitions and Properties*

Suppose there exists a random variable X . The definition of probability in the continuous case presumes that there exists a function f , called the probability density function, such that areas under the curve give the probabilities associated with the corresponding intervals along the horizontal axis. In order for a function to be a density function, the area under the graph must be equal to one, and the function is always greater to or equal to zero. Therefore, one should keep this in mind when trying to fit a curve to the data.

A density function gives a natural description of the distribution of the random variable. In this chapter, the following problem is addressed: Given a set of data points from a sample of an unknown probability density function, can one construct an estimate of the density function?

Density estimation can either be parametric or non-parametric. In the parametric method, one assumes that the observed data points are taken from a distribution already known like for example the normal distribution. The work of the statistician will then be to estimate the mean and variance of the distribution. Parametric estimation is a bit rigid over the observed data, while non-parametric estimation does not give the value of the fitted curve, but only the probabilities.

But does one really need density estimation? Density estimation can be used as a way of displaying the data. It can give indication of the type of distribution that the sample under consideration has. For example, it can show skewness, multimodality,

etc. In fact the basic histogram is in itself a very crude density estimation and will be considered in the next section.

The following diagram gives an example of kernel estimation of the expenditure from the observed data of the Household budget survey. As one can see from the graph, the majority of people spend around Lm5000 a year. Also, one can see the expenditure is unimodal.

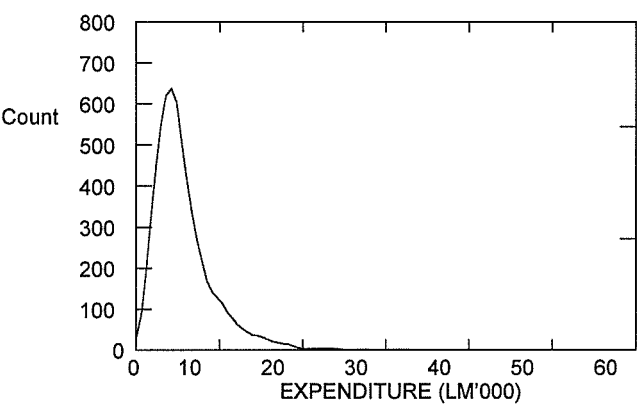


Figure 3.1: Kernel estimation of expenditure

Density estimation is becoming popular because the estimates are comprehensible to non-mathematicians. This is because one can easily look at the above plot and read certain aspects of the figure, such as one mode, etc.,

In the following sections, the main methods available for density estimation will be considered. Although density will be estimated for univariate cases only, one must keep in mind that many of the important applications of density estimations are to multivariate data. The total expenditure data points (from the HBS) will be considered as examples throughout these sections.

3.2 *Measurements of Error*

On estimating the densities, one must keep in mind that the density obtained is never as good as the data. Thus, before going into density estimation, some measurements of error must be introduced. Let f be the true density of the function, and $\hat{f}$ be the density estimator. The reader can easily acknowledge the fact that it is impossible that f equals $\hat{f}$ at all points. Thus, one would need to consider the error between the original function and the estimated one. When considering estimation at a single point x , a natural measure is the **Mean Square Error (MSE)**.

The MSE is defined as

$$MSE_x(\hat{f}) = E\left(\hat{f}(x) - f(x)\right)^2 \dots\dots\dots(3.1)$$

The above expression is simply subtracting the actual point from the estimated point, and then averaging all over the sample.

By opening up the expression, one gets

$$MSE_x(\hat{f}) = E\left(\hat{f}(x)\right)^2 + E\left(f(x)\right)^2 - 2E\left(\hat{f}(x)\right)f(x)$$

Consider the expression

$$E\left(\hat{f}(x) - f(x)\right)^2 + E\left(\hat{f}(x) - Ef(x)\right)^2$$

On factorising this expression, one gets

$$E\hat{f}(x)E\hat{f}(x) - 2E\hat{f}(x)f(x) + f(x)f(x) + E\left(\hat{f}(x)\hat{f}(x)\right) - 2E\hat{f}(x)Ef(x) + Ef(x)Ef(x)$$

which simplifies to

$$E\left(\hat{f}(x)\right)^2 + E(f(x))^2 - 2E\left(\hat{f}(x)\right)f(x)$$

Besides this, one can note that $E\left(\hat{f}(x) - f(x)\right)^2 + E\left(\hat{f}(x) - Ef(x)\right)^2$ is in fact the squared bias and the variance at the point x . Thus the bias can be reduced at the expense of increasing the variance and vice versa.

The **Mean Integrated Square Error (MISE)** is a measure on the global accuracy of $\hat{f}$ as an estimator of f . It is defined by

$$\text{MISE}\left(\hat{f}\right) = E\left[\int\left(\hat{f}(x) - f(x)\right)^2 dx\right]$$

One can note that this equation is similar to the formula for MSE, and the integration ensures that this measurement is global and not pointwise.

Using equation (3.1) then

$$\text{MISE}\left(\hat{f}\right) = \int \text{MSE}_x\left(\hat{f}\right) dx = \int\left(E\hat{f}(x) - f(x)\right)^2 dx + \int \text{var } \hat{f}(x) dx$$

That is MISE is the sum of the integrated square bias and the integrated variance.

3.3 *Non-Parametric Estimation*

3.3.1 *The Histogram*

By definition a histogram is a graphical representation of a frequency density. It consists of rectangles known as bins, having their bases on the horizontal x -axis with centres at the class marks and lengths equal to the frequency of that class interval. For each bin, the area is proportional to the class frequency. A frequency histogram can be transformed into a density histogram by normalising such that the histogram integrates to one. Given an origin x_0 and a bin width h , the bins of the histogram are

defined to be the intervals $[x_0 + mh, x_0 + (m+1)h)$ for positive and negative integers of m .

Thus, the histogram is defined as

$$\hat{f}(x) = \frac{1}{nh} (\text{no. of } X_i \text{ in same bin as } x)$$

where $X_1, \dots, X_n$ denote observations

$\hat{f}$ is the density estimator.

However, in this way, one has to define both the origin and the bin width. Consider the following two diagrams.

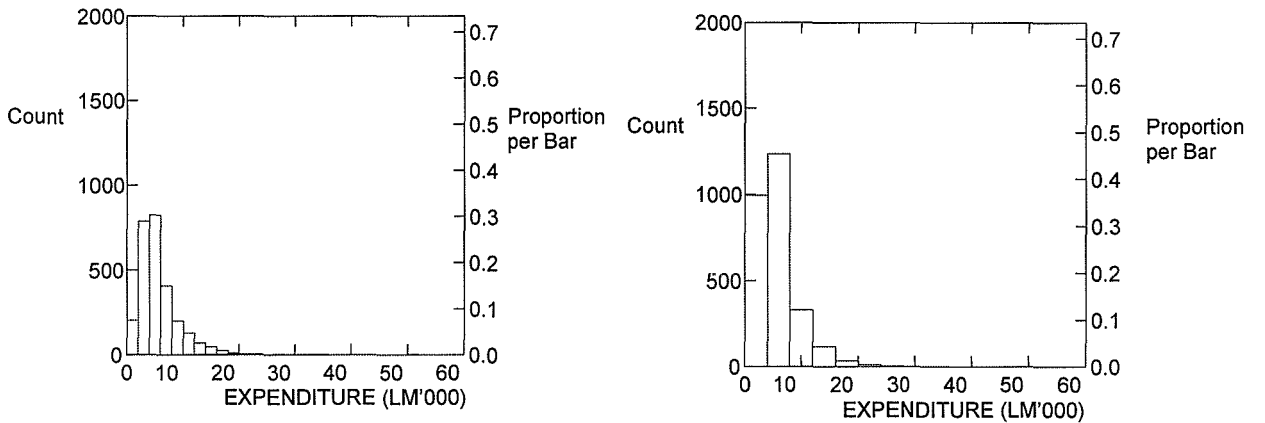


Figure 3.2 Histogram of expenditure with different bin widths and different number of bins

They both depict the consumer expenditure across the year, however they have different bin widths and number of bins. The first diagram has 42 bins, while the second diagram has 15 bins. One can see what a great difference this makes to the shape of the histogram. Besides the difference in shape, a major disadvantage in the histogram is that it is discontinuous, thus creating problems when derivatives of estimates are required.

Although the origin of the histogram goes back to 1662, it was Herbert Sturges³⁴ in 1926 who gave systematic guidelines for designing histograms. He suggested taking the Normal density as a point of reference when thinking about the data. Besides, he observed that the binomial distribution $B(n, p = 0.5)$ could be used as a model of an optimally constructed histogram with appropriately scaled Normal data. The algorithm for constructing the histogram is as follows:

1. Construct a frequency polygon with k bins, each of width 1 and centred on the points $i = 0, 1, \dots, k-1$.

2. Choose the bin count of the i^{th} bin to be the Binomial coefficient $\binom{k-1}{i}$. As k

increases, this ideal frequency histogram assumes the shape of a normal density

with mean $\left(\frac{k-1}{2}\right)$ and variance $\left(\frac{k-1}{4}\right)$. The total sample size is

$$n = \sum_{i=0}^{k-1} \binom{k-1}{i} = \binom{k-1}{0} + \binom{k-1}{1} + \dots + \binom{k-1}{k-1}$$

$$\frac{(k-1)!}{0!(k-1)!} + \frac{(k-1)!}{1!(k-2)!} + \dots + \frac{(k-1)!}{(k-1)!0!}$$

which is the binomial expansion of $(1+1)^{k-1} = 2^{k-1} = n$.

Therefore

$$(k-1) = \log_2 n \Rightarrow k = 1 + \log_2 n$$

This is a rule for the number of bins that a histogram must have. Thus, if the household budget survey contains 2722 data points, substituting this for n , one will get the ideal number of bins for a histogram. Therefore $k = 1 + \log_2 2722 = 12.4$.

Thus the ideal number of bins will be 12 or 13, and dividing the sample range of the

data into the prescribed number of equal width bins. However, if the data are not normal (as in most cases), additional bins may be required. In fact, Doane suggested that the number of bins is increased by $\log_2 \left(1 + \hat{\gamma} \sqrt{n/6} \right)$, where $\hat{\gamma}$ is an estimate of the standardised skewness coefficient. Terrel and Scott (1985) showed that the minimum number of bins is $\sqrt[3]{2n}$, so in the case of the HBS, one should need $\sqrt[3]{5444} = 18$ bins.

The following is a histogram of the total expenditure using 15 bins and 18 bins.

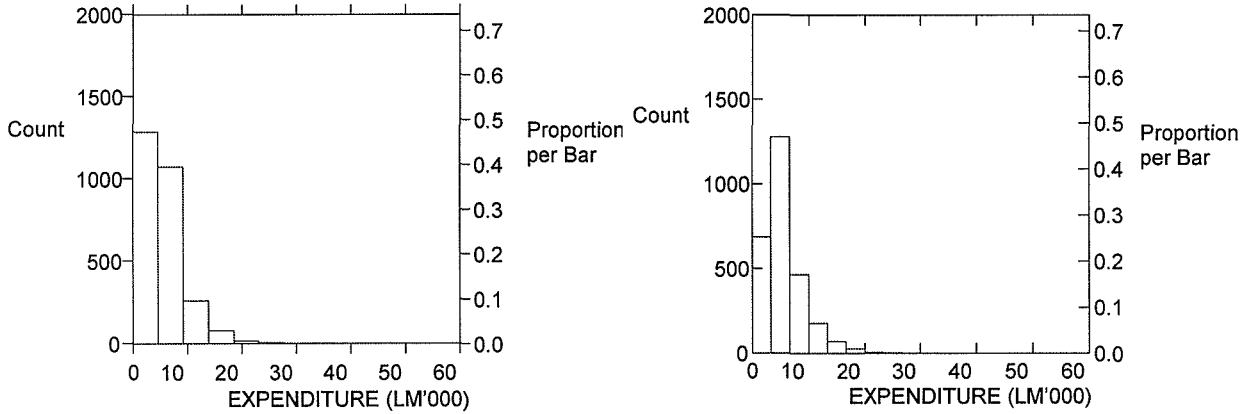


Figure 3.3 Histogram of expenditure with 15 and 18 bins respectively

In order to find the MSE and MISE of the histogram, a density histogram with building blocks each of $\frac{1}{nh}$ is considered, where h is the fixed bin width of the histogram. Let the number of points falling in the k^{th} bin (B_k) be denoted by ν_k . Thus the density histogram is defined as

$$\hat{f}(x) = \frac{\nu_k}{nh} = \frac{1}{nh} \sum_{i=1}^n I_{[t_k, t_{k+1}]}(x_i)$$

for $x \in B_k$

Since the bin counts are Binomial random variables, i.e.

$$\nu_k \sim B(n, p_k)$$

where $p_k = \int_{B_k} f(t) dt$, then using Binomial rules

$$E[\nu_k] = np_k$$

and

$$\text{Var}[\nu_k] = np_k(1 - p_k)$$

Therefore

$$\text{Var} \hat{f}(x) = \frac{\text{Var}(\nu_k)}{(nh)^2} = \frac{p_k(1 - p_k)}{nh^2}$$

while

$$\text{Bias} \hat{f}(x) = E \hat{f}(x) - f(x) = \frac{1}{nh} E(\nu_k) - f(x) = \frac{np_k}{nh} - f(x) = \frac{p_k}{h} - f(x)$$

3.3.2 The Naive Estimator

The naive estimator tries to improve the histogram such that it is independent of the choice of the bin positions. To construct this estimation, one needs to make use of the following function

$$\hat{f}(x) = \frac{1}{n} \sum_{i=1}^n \frac{1}{h} w\left(\frac{x - X_i}{h}\right)$$

where w is defined as $w(x) = \begin{cases} \frac{1}{2} & \text{if } |x| < 1 \\ 0 & \text{otherwise} \end{cases}$

However, one can see that the choice of the bin width h still remains, and it is this, which controls the amount by which data are smoothed. It follows that the naïve

estimator will be just an average of the data points falling inside a particular bin width. However this density estimate is not a continuous function since it has jumps at every $X_i \pm h$, and it has zero derivative elsewhere. The naïve estimator was used on the total expenditure of the HBS. First, one has to determine the value of the bin width h . Ignoring the last value of the total expenditure, which is a very high value relative to the others (hence it is an outlier), the bin range is divided such that the number of bins is equal to 19. Thus h will be 1956.15. The frequency of each bin was subsequently found. Then, a program was constructed using Pascal (refer to Appendix A) such that it automatically calculates the value of w , whether it has a value of 0 or a value of 0.5. After this procedure, each value was divided by h , and summed over each data point. Finally the values were divided by the value of n for normalising. The estimate thus constructed yields the following graph.

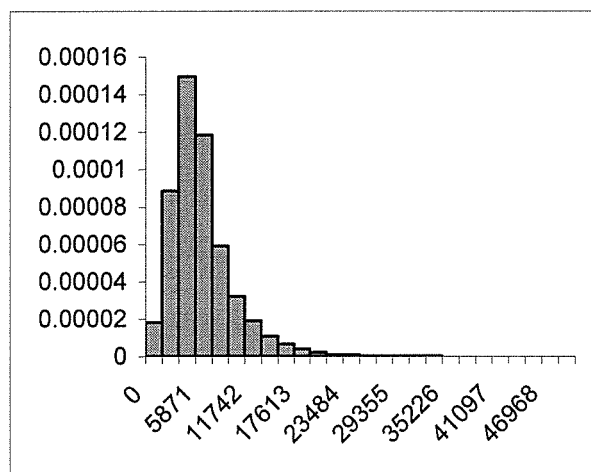


Figure 3.4 *Histogram after Naive estimation of expenditure*

On checking the area, one can find that it is in fact equal to 1.

3.3.3 The Kernel Estimator

Basically, the kernel estimator is a generalisation of the naïve estimator since it uses the same equation

$$\hat{f}(x) = \frac{1}{n} \sum_{i=1}^n \frac{1}{h} w\left(\frac{x - X_i}{h}\right)$$

but in this case, w is replaced by a function K called the kernel, satisfying the

condition that $\int_{-\infty}^{\infty} K(x) dx = 1$. Thus the kernel estimator can be written as

$$\hat{f}(x) = \frac{1}{nh} \sum_{i=1}^n K\left(\frac{x - X_i}{h}\right)$$

or, as suggested by Rosenblatt (1956),

$$\hat{f}(x) = \frac{1}{n} \sum_{i=1}^n K_h(x - X_i)$$

A common function used for K is the normal density function. In this case, h is called the window width. Instead of constructing boxes and estimating a curve over the boxes, as in previous cases, the kernel estimation involves estimating a curve from other curves, or as some of the authors refer to it, it is a sum of the ‘bumps’ placed at the observations. The shape of the bumps is determined by the kernel function, while h determines their width. h also determines the level of detail estimated. In fact as $h \longrightarrow 0$, the estimate will have a spike at every observation, whereas as h becomes large, all detail will be obscured.

In many cases, K is chosen to be a non-negative function, therefore it would be a probability density function. Hence, it follows that $\hat{f}(x)$ will itself be a probability density. Furthermore, it will inherit all the continuity and differentiability properties of the kernel.

Although kernel estimation is one of the major estimation methods used nowadays, it suffers from a slight drawback when applied to long-tailed distributions. Since h must be fixed throughout the sample, there is a tendency for false noise to appear in the tails of the estimates. If the tails are smoothed, this will over-smooth the whole sample, thus leaving out possible features of the data.

Parzen³⁴ has shown that on minimising the MISE of the kernel estimation, given by

$$MISE = \frac{1}{4} h^4 k_2^2 \int f''(x)^2 dx + n^{-1} h^{-1} \int K(t)^2 dt$$

the ideal window width is given by

$$h_{opt} = k_2^{-2/5} \left\{ \int K(t)^2 dt \right\}^{1/5} \left\{ \int f''(x)^2 dx \right\}^{-1/5} n^{-1/5}$$

where

$$k_2 = \int t^2 K(t) dt$$

This optimum window depends on the density being estimated. However, on substituting the value obtained for h into the formula for MISE, then the approximate value of the mean integrated square error will be

$$\frac{5}{4} C(K) \left\{ \int f''(x)^2 dx \right\}^{1/5} n^{-4/5}$$

where
$$C(K) = k_2^{2/5} \left\{ \int K(t)^2 dt \right\}^{4/5}$$

The MISE will be a small number if the kernel is chosen in such a way such that $C(K)$ is also small. The problem is now therefore that of minimising $C(K)$. k_2 need not be equal to one, because if it is not, one can replace the kernel by $k_2^{1/2} K(k_2^{1/2} t)$, which will not effect the value of $C(K)$. Therefore the problem reduces to that of minimising $\int K(t)^2 dt$ subject to $\int K(t) dt = \int t^2 K(t) dt = 1$. Hodges and Lehmann³⁴ solved the problem by setting $K(t)$ to be the function:

$$K(t) = \frac{3}{4\sqrt{5}} \left(1 - \frac{1}{5}t^2 \right)$$

when $-\sqrt{5} \leq t \leq \sqrt{5}$, and 0 otherwise.

This is called the Epanechnikov kernel because it was suggested by Epanechnikov³⁴ (1969). He even suggested that the choice of kernel is not crucial to the statistical performance of the method, and therefore it is quite reasonable to choose a kernel for computational efficiency. In fact Epanechnikov kernel is also locally optimal for estimating a regression function.

Other kernels suggested in various texts are the Gaussian kernel³⁵, i.e.

$$K(t) = \frac{1}{\sqrt{2\pi}} e^{-(1/2)t^2}$$

and the Rectangular kernel³⁷, i.e.

$$\frac{1}{2} \text{ for } |t| < 1, \text{ and } 0 \text{ otherwise}$$

If the Gaussian kernel is used, $k_2 = 1$ and $h_{opt} = C_0 n^{-1/5}$, where

$$C_0 = (4\pi)^{-1/10} \left\{ \int f''(t)^2 dt \right\}^{-1/5}$$

The kernel method is a useful technique in analysing data. It provides an attractive method for estimation of both probability densities and their derivatives. It need not assume a parametric form for the detection function, thus making it a data-oriented modelling approach. However one must choose the smoothing width h , which controls the smoothing of the curves fitted. A problem pointed out by Buckland³⁵ is that the kernel method cannot be applied to data subjected to severe rounding error. This is because this will result in a smaller width h and an inflated estimate for $f(0)$.

3.4 *Parametric Density Estimation*

In this previous sections, major non parametric ways of how a curve can be estimated was introduced. However, these methods do not estimate the parameters, they only estimate a curve through the points. Therefore to obtain some parameters, one needs to look into parametric density estimation.

3.4.1 *The Least Squares Method*

The first method tried on the expenditure of the HBS was the Least Squares Method. This method requires that one minimises the sum of squares of the deviation of the point from a specified curve. The user must choose the curve, and in many cases the curve chosen is not the correct one.

One approach applied to the data was the Cubic function approach.

The function used was

$$f(x) = ax^3 + bx^2 + cx + d$$

Frequencies were converted to probabilities by dividing by the total number of points considered.

Integrating

$$f(x) = ax^3 + bx^2 + cx + d$$

one gets

$$f(x) = \frac{ax^4}{4} + \frac{bx^3}{3} + \frac{cx^2}{2} + dx$$

Then integrating interval by interval, one gets

$$ax_i + by_i + cz_i + du_i$$

The difference from the real values of areas was then found, i.e.

$$ax_i + by_i + cz_i + du_i - p_i$$

This difference needs to be squared and summed up to give

$$\sum_i (ax_i + by_i + cz_i + du_i - p_i)^2$$

Differentiating with respect to a , b , c and d respectively is carried out in order to minimise the squared difference, i.e.

$$\frac{\partial}{\partial a} = \frac{\partial}{\partial b} = \frac{\partial}{\partial c} = \frac{\partial}{\partial d} = 0$$

This will yield the following equations:

1. $a \sum x_i^2 + b \sum x_i y_i + c \sum x_i z_i + d \sum x_i u_i = \sum x_i p_i$
2. $a \sum x_i y_i + b \sum y_i^2 + c \sum y_i z_i + d \sum y_i u_i = \sum y_i p_i$
3. $a \sum x_i z_i + b \sum y_i z_i + c \sum z_i^2 + d \sum z_i u_i = \sum z_i p_i$
4. $a \sum x_i u_i + b \sum y_i u_i + c \sum z_i u_i + d \sum u_i^2 = \sum u_i p_i$

The above equations are then transformed into a matrix, from which values of a , b , c and d can be found.

3.5 *Curve Fitting using Parametric Estimation*

Since the objective of this study was to get a final curve that fitted the data, one can obviously see that non-parametric estimation is not suitable for this study. Both the least-squares estimation and the maximum likelihood estimation were tried on the data. However the least squares estimation was inappropriate because in certain cases, equation became too complicated and impossible to solve unless one goes into more complex methods, which was not the objective of this study. Maximum Likelihood Estimation was inappropriate because the errors obtained were not normally

distributed, which is a condition that the data must satisfy if one is going to use this method.

Since no other method was found suitable for finding a curve that fits the data points, it was decided that the data points should be transformed into a histogram. This helped in determining the shape of the distribution. A Qbasic program was constructed (check Appendices A for the programs), in which the best-fit curve was determined empirically. Finally, a goodness of fit test was applied to check whether the density obtained should be rejected or not.

This method was repeated for all the commodities in the HBS together with the total expenditure. What follows is the results for all the commodities.

3.5.1 *Expenditure*

Expenditure is found from the HBS by adding the expenditure of all the commodities. The data points for the total expenditure were grouped into a histogram. The histogram obtained using Excel is given in the following diagram.

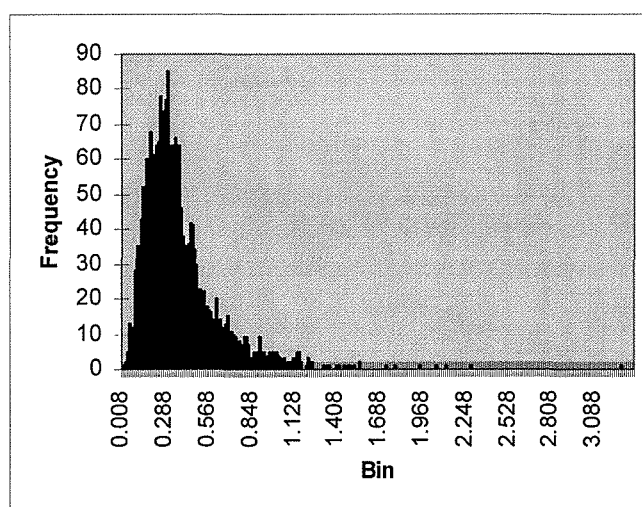


Figure 3.5 *Histogram of Total Expenditure (units $\times 10^3$)*

A total of 156 bins were considered for the histogram. On looking at the histogram, one could see that it looks like a lognormal distribution. Hence the data points were imported into SPSS, and a p-plot was applied, testing the hypothesis that it is lognormally distributed.

The P-Plot is a facility available in SPSS that plots the observed with the expected cumulative probability. This is indicated by the thin line. The thick line gives the plot of the real data. If the two appear over each other, then that is the best curve you can get. On looking at the p-plot, one could see that the data points do not fit a gamma distribution exactly because in the first part there is under estimation of the data, and in the second part there is over estimation.

The p plot obtained was as follows:

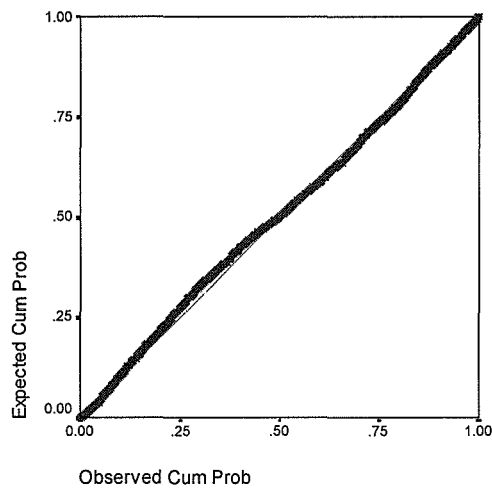


Figure 3.6 **Lognormal P-Plot of Total Expenditure**

This shows that the distribution is in fact a lognormal distribution. The lognormal curve fitted had $m = 0.34$ and $\sigma = 0.53$. The following diagram shows the two graphs superimposed on each other.

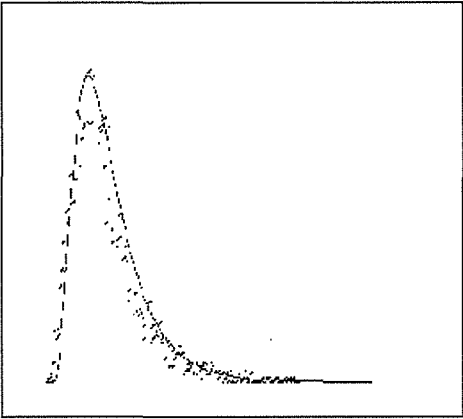


Figure 3.7 *Lognormal Curve Fitted to Total Expenditure*

Thus the equation for the total expenditure is given by

$$f(x) = \frac{1}{x(0.53)(2\pi)^{1/2}} \exp \left\{ \frac{-[\log(x/(0.34))]^2}{2(0.53)^2} \right\}$$

The test applied for goodness of fit was the Kolmogorov-Smirnov test. This investigates the difference between an observed distributions and a specified population distribution. This density estimate was accepted at a level of 0.01 significance. The following table gives some descriptive statistics of the total expenditure.

	Minimum	Maximum	Mean		Std. Deviation	Variance
	Statistic	Statistic	Statistic	Std. Error	Statistic	Statistic
Expenditure	.00	.0313	6.410E-03	6.761E-04	8.444E-03	7.131E-05

Table 3.1 *Descriptive Statistics of Total Expenditure*

3.5.2 Food

The next distribution considered was the distribution of the budget shares of ‘Food’. Food budget shares are obtained by dividing the original values of food expenditure by the total expenditure.

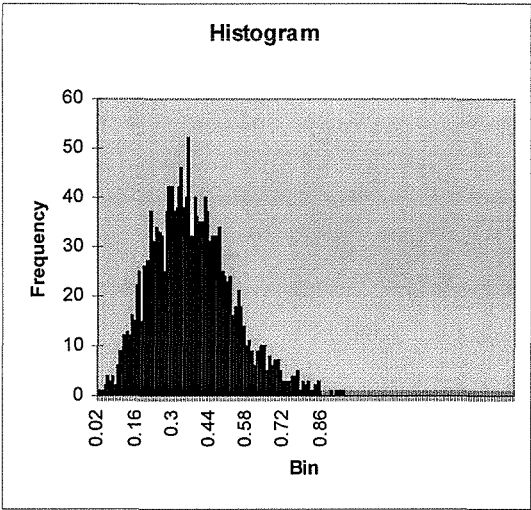


Figure 3.8 *Histogram of Food Budget Share(*10³)*

This looks like a Weibull distribution. In fact the following p-plot was obtained.

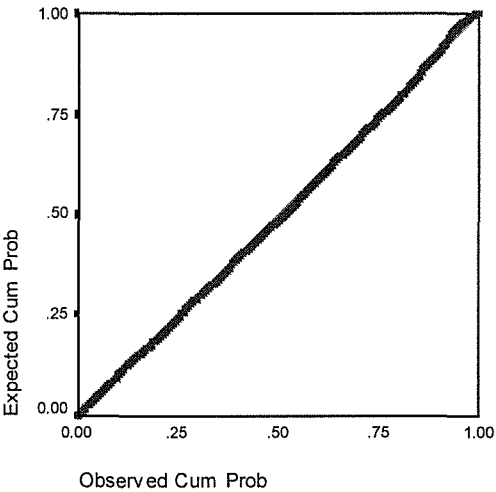


Figure 3.9 *Weibull P-Plot of Food*

The above figure shows that trying to fit a Weibull distribution would be a good idea. The parameters that best fitted the data points were $b = 0.68$, $c = 2.7$ and $s = 0.09$, where s is the shift in the x-axis. The two graphs superimposed on each other are given in the following diagram.

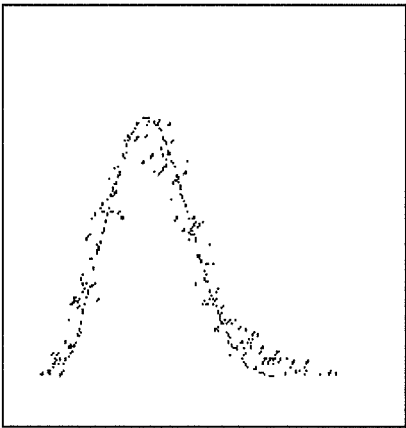


Figure 3.10 *Weibull Curve Fitted to Food*

The equation fitted is given by

$$f(x + 0.09) = 15 \left(\frac{2.7x^{1.7}}{0.68^{2.7}} \right) \exp \left[- \left(\frac{x}{0.68} \right)^{2.7} \right]$$

The Kolmogorov-Smirnov test was once again accepted with a level of 0.01 significance. The descriptive statistics are as follows:

	Minimum	Maximum	Mean		Std. Deviation	Variance
	Statistic	Statistic	Statistic	Std. Error	Statistic	Statistic
Food	.0000	.0202	6.30428E-03	4.37815E-04	5.46830E-03	2.990E-05

Table 3.2 *Descriptive Statistics of Food Budget Share*

Some outliers removed previously from the data were then analysed. The outliers are three couples having an average income of approximately Lm2500. All three spend more than half of their income on food. All the households were living on a pension. However, as one can see from the raw data, the average expenditure of households on food tends to increase with age.

3.5.3 Beverages and Tobacco

The second commodity considered was ‘Beverage and Tobacco’. Although the HBS considers household expenditure of this commodity, one must take into account the fact that in most cases, only the parents spend money for alcohol and tobacco. The histogram obtained is shown in the following diagram.

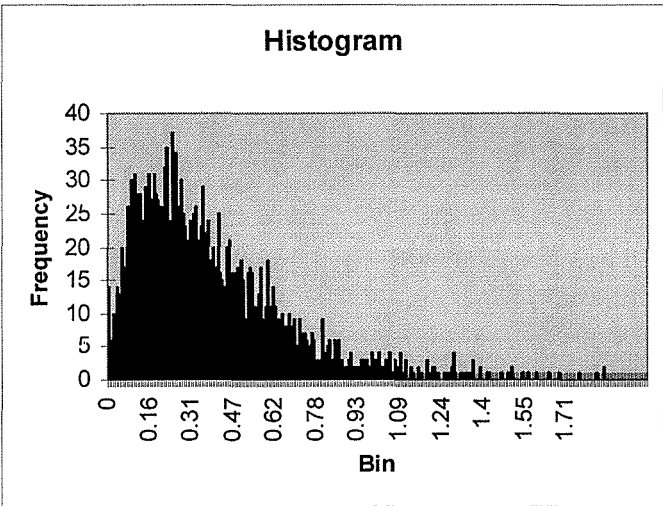


Figure 3.11 Histogram of Beverage and Tobacco Budget Share ($\ast 10^{-3}$)

This looks like a gamma distribution. On testing the hypothesis that the data points are gamma distributed, the following p-plot was obtained.

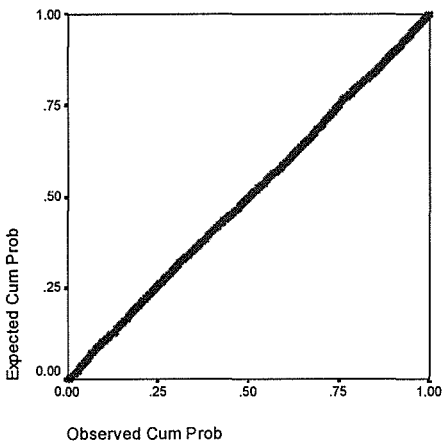


Figure 3.12 Gamma P-Plot of Beverage and Tobacco

The gamma curve obtained is shown in the following diagram.

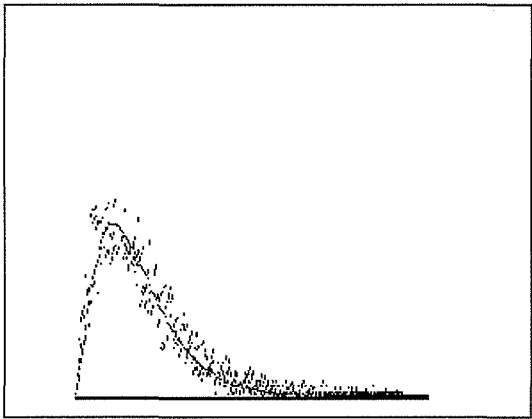


Figure 3.14 *Gamma Curve Fitted to Beverage and Tobacco*

The data points were multiplied by a factor of 40 to fit the screen. The parameters for the gamma distribution were $c = 2.5$ and $b = 0.28$.

Thus the equation is given by

$$f(x) = 2\left(\frac{x}{0.28}\right)^{1.5} \frac{\exp\left(\frac{-x}{0.28}\right)}{0.28\Gamma(2.5)}$$

where $\Gamma(2.5)$ is the gamma function with parameter 2.5, given by

$$\Gamma(2.5) = \int_0^{\infty} \exp(-u)u^{1.5}du$$

Note that $\Gamma(0.5) = \sqrt{\pi}$.

The descriptive statistics are given by

	Minimum	Maximum	Mean		Std. Deviation	Variance
	Statistic	Statistic	Statistic	Std. Error	Statistic	Statistic
Food	.0000	.0136	2.70270E-03	1.78123E-04	3.42626E-03	1.174E-05

Table 3.3 *Descriptive Statistics of Beverage and Tobacco Budget Share*

If the outliers are identified, one notices that one of the outliers is a household an old woman who lives alone. Besides having a pension, she has a part-time job and also some investments, which in total amount to 2016.88. Looking at her expenditures, one can see that the major part of it is spent on food, and on beverages and tobacco. One could argue that maybe she is an alcoholic or a chain smoker, however, since such data is not available, one can only make suppositions. The second outlier is a young family whose average age is 28. The man is a machine operator with an income of around Lm3000. Their major expenditure is on beverages and tobacco and then on food. Their expenditure on housing is minimal, and they spend nothing on health.

3.5.4 *Clothing and Footwear*

The same procedure was repeated for the thirdrd commodity ‘Clothing and Footwear’. The histogram obtained is shown in the following diagram.

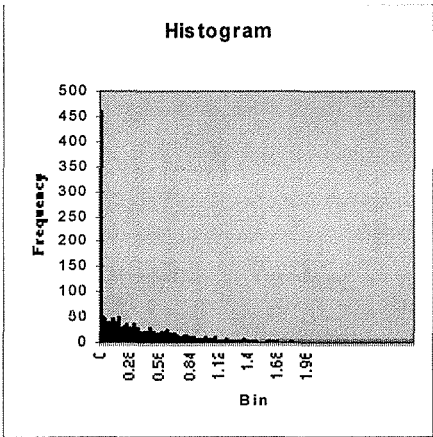


Figure 3.14 *Histogram of Clothing and Footwear Budget Share(*10⁻³)*

One can see that a large number of the sample considered (in fact 460 household) spend nothing on this commodity. Besides this, except for the first part of the graph, the distribution is an exponential distribution, as can be confirmed by the p-plot.

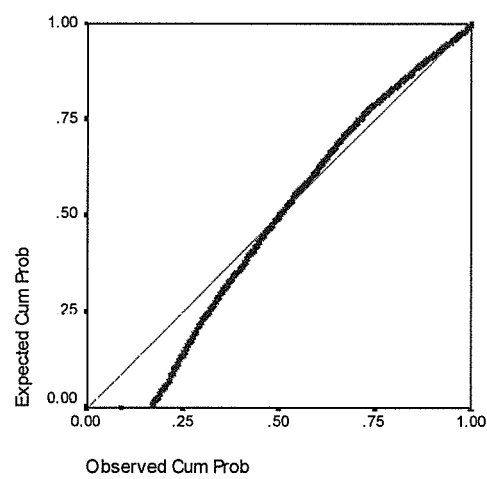


Figure 3.15 *Exponential P-Plot of Clothing and Footwear*

In this case, the Qbasic generated exponential data points such that it fits a curve through the original data points. The data points were multiplied by a value of 70 to fit the screen. The fitted curve is shown in the following diagram.

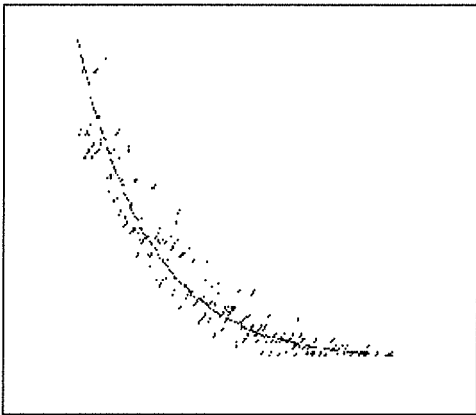


Figure 3.16 *Exponential Curve Fitted to Clothing and Footwear*

The parameters used where $\lambda = 2.2$ and $s = -0.22$, where s is the shift in the x-axis. Therefore, the equation is given

$$f(x - 0.22) = \frac{2.2}{70} \exp(-2.2x)$$

The descriptive statistics are given by

	Minimum	Maximum	Mean		Std. Deviation	Variance
	Statistic	Statistic	Statistic	Std. Error	Statistic	Statistic
Clothing & Footwear	.0000	.1695	5.0761E-03	8.9954E-04	1.2625E-02	1.594E-04

Table 3.4 ***Descriptive Statistics of Beverage and Tobacco Budget Share***

On looking at the outliers from the data, one can identify a particular household made up of an old woman living on a pension and other social income of around Lm1800. Strangely enough, she spends more than Lm2000 on clothes and footwear. One could also see that many young households spend nothing on clothing and footwear. This could be due to the fact that these couples have house loans to pay for, and so they spend money only on necessities, and clothes and footwear do not fall under this category. Even households made up of old people rarely spend anything on this commodity, as expected. Therefore, by looking at the data, one can see that many households consider this commodity as a luxury.

On applying the goodness of fit test, the density estimated was rejected. This may be due to the fact that the first part of the curve was not an exponential curve.

3.5.5 *Housing*

One would think that this commodity depends on the age of the household considered, because an old household would have paid all the loans (if any) on the house, while young households would still have a bulk to pay. Besides that, other households tend to spend nothing or very little on housing. The histogram obtained is as follows.

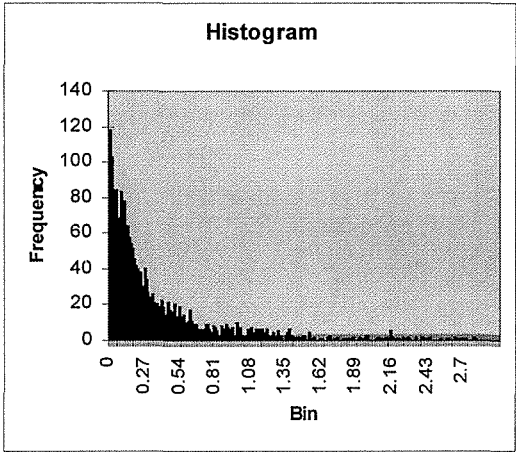


Figure 3.17 *Histogram of Housing Budget Share (*10⁻³)*

As can be seen from the graph, the majority of people spend nothing or very little for the home. The curve looks like an exponential distribution. On applying the p-plot, one can see that the exponential distribution could in fact fit the data.

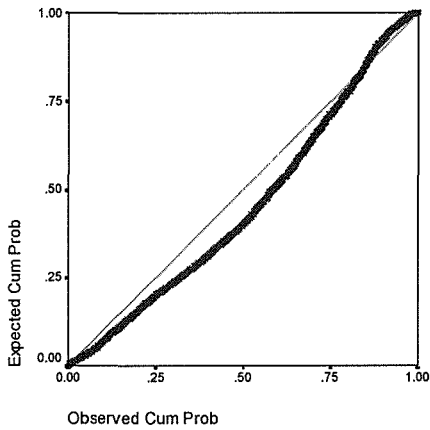


Figure 3.18 *Exponential P-Plot of Housing*

However, the goodness of fit test rejected the density estimated because the exponential distribution is over estimating the data.

The curve fitted is shown in the next diagram.

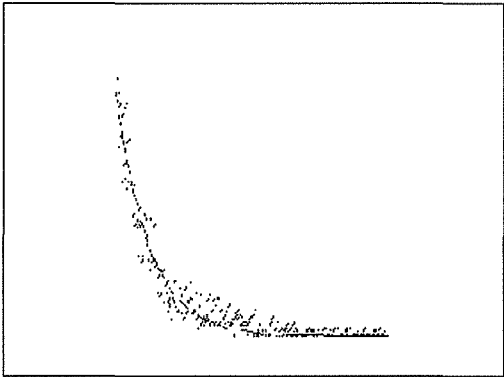


Figure 3.19 *Exponential Curve Fitted to Housing*

The final result of the exponential was shifted in the x-axis by a value of -0.15, and the other parameter $\lambda = 3.2$.

Thus the equation fitted is given by

$$f(x - 0.15) = 3.2 \exp(- 3.2x)$$

The descriptive statistics are given by

	Minimum	Maximum	Mean		Std. Deviation	Variance
	Statistic	Statistic	Statistic	Std. Error	Statistic	Statistic
Housing	.0000	.0435	3.5587E-03	4.1757E-04	6.9998E-03	4.900E-05

Table 3.5 *Descriptive Statistics of Housing Budget Share*

Looking at the outliers, one can see that the largest outlier is a family both over 68 years of age. They have an income of around Lm3000, from pension and social services. Their largest expenditure is on housing; in fact they spend more than Lm3600. The couple could either have a rent of Lm300 per month. Another

possibility is that the couple is living in a centre for old people, either governmental or private. This could be an explanation for the high rent.

Two other outliers are two old women, over 64 years of age. One of the women is a pensioner with income of around Lm2700. Her major expense is on housing, and then on food. Besides having a pension, the other woman also has some investments and a part time job with a total income of around Lm6000. She spends more than Lm3000 on housing. One can think that she also lives in a house for old people, thus paying a high rent. However, this woman also spent almost Lm3000 on furniture and equipment for the house. It could be that this woman changed or bought a new house during the survey year, and that would explain her high expenditure on house related commodities.

Looking at households having young people living in it, one can see that many young households spend little or nothing on housing, however there are some households, which have a relatively large budget share. These could be the households paying the loans for their homes or those households that have just bought a house. In fact, there is a family who has an expenditure of more than Lm12000 on housing. Another family with an income of around Lm2000 has an expenditure of more than Lm4000 on housing.

3.5.6 *Fuel and Power*

Fuel and power is being considered as a necessity for many people. The majority of people now have a car, and in fact many households possess more than 1 car. The histogram for this distribution is as follows.

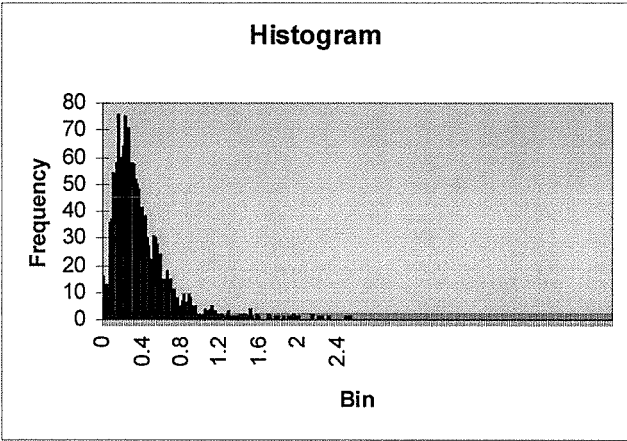


Figure 3.20 *Histogram of Fuel and Power Budget Share(*10⁻³)*

One can see that a lognormal distribution could in fact fit the data. So, a P-plot to test the shape of the distribution was applied, and the following diagram obtained.

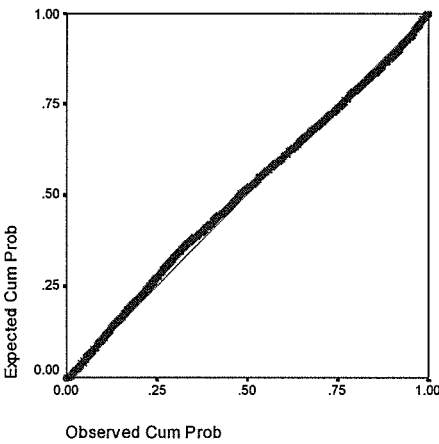


Figure 3.21 *Lognormal P-Plot of Fuel and Power*

However, one must note that in order to run the lognormal p-plot, data points with a value of 0 had to be ignored, as the lognormal distribution is only defined for positive values. The following density curve was estimated.

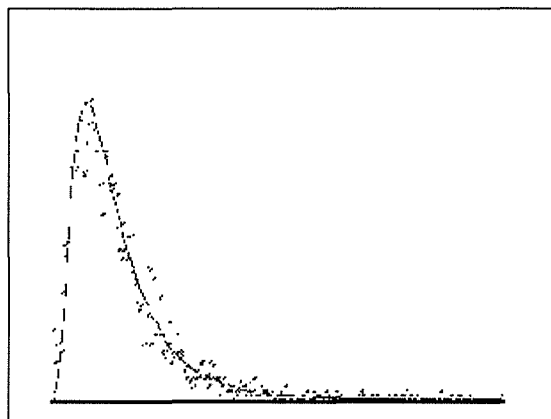


Figure 3.22 *Lognormal Curve Fitted to Fuel and Power*

The parameters for the distribution are $m = 0.3$, $\sigma = 0.65$, $s = 0$ and $d = 1.1$, where s is the shift in the x-axis, and d is the parameter that scales the curve down. Thus, the equation is given to be

$$f(x) = \frac{1.1}{x(0.65)(2\pi)^{1/2}} \exp \left\{ \frac{-[\log(x/(0.3))]^2}{2(0.65)^2} \right\}$$

The descriptive statistics are given by

	Minimum	Maximum	Mean		Std. Deviation	Variance
	Statistic	Statistic	Statistic	Std. Error	Statistic	Statistic
Fuel & Power	.0000	.0280	3.9370E-03	4.2524E-04	6.7772E-03	4.593E-05

Table 3.6 *Descriptive Statistics of Fuel & Power Budget Share*

One of the outliers is a household with just an old woman having an income of about Lm1700. The woman spends mostly on food, then on health, and then on fuel and power in that order. She spends nothing on the other commodities. The second outlier is an old family with average age of 81. Their income is around Lm6000. They spend almost Lm850 on food, and more than Lm300 on fuel and power.

Another outlier is a young couple with two small children, in fact average age of family is seventeen. They spend mostly on fuel and power, more than Lm900, and their income is around Lm4000.

The goodness of fit test accepted the density curve with 0.05 level of significance.

3.5.7 Furniture and Household Equipment

One would think that this commodity, like others mentioned before, is affected by age, since young people starting their new families tend to spend more on their homes than old people. The same procedure for getting the budget shares and hence a histogram was applied. The resulting histogram is given in the following diagram.

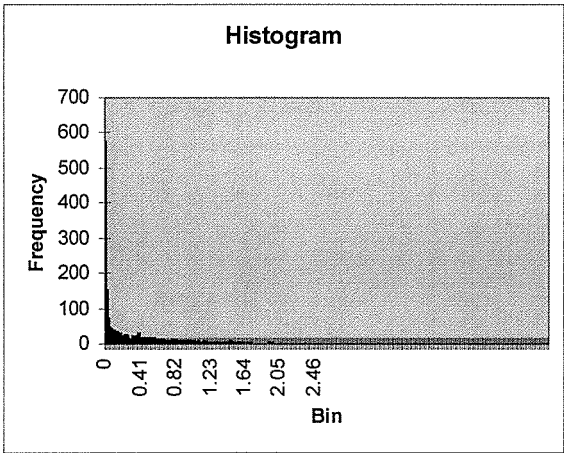


Figure 3.23 Histogram of Furniture and Household Equipment Budget Share (*10⁻³)

One can see that in fact more than 500 households spend nothing on furniture and household equipment. So, one can presume that this commodity is considered as a luxury commodity throughout the households of Malta. Besides, old households spend almost nothing on this commodity.

Although no distribution was found to be exactly the same as the data points, the gamma distribution was found to be the best guess for the curve to the data points, as seen by the following diagram.

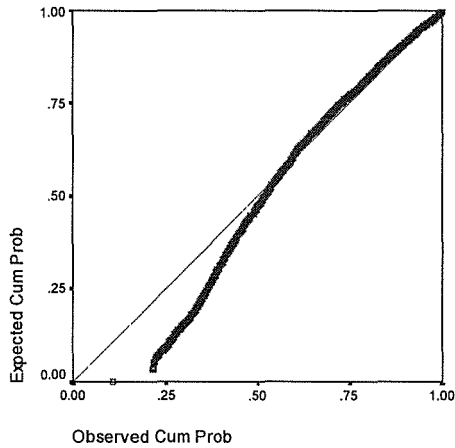


Figure 3.24 ***Gamma P-Plot of Furniture and Household Equipment***

The following curve was obtained.

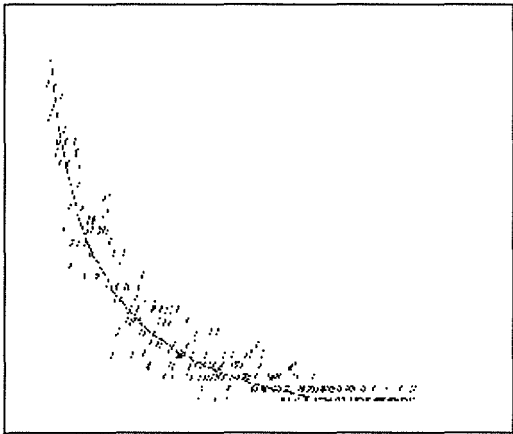


Figure 3.25 ***Gamma Curve Fitted to Furniture and Household Equipment***

The parameters used are $c = 0.5$, $b = 1.5$ and $d = -0.07$, where d is the shifting parameter.

The equation is given

$$f(x) = 0.07 + \left(\frac{x}{1.5}\right)^{0.5} \frac{\left[\exp\left(\frac{-x}{1.5}\right)\right]}{1.5\Gamma(0.5)}$$

The descriptive statistics are given by

	Minimum	Maximum	Mean		Std. Deviation	Variance
	Statistic	Statistic	Statistic	Std. Error	Statistic	Statistic
Furniture & Household Equipment	.0000	.2127	3.5087E-03	8.0106E-04	1.3523E-02	1.829E-04

Table 3.7 ***Descriptive Statistics of Furniture & Household Equipment Budget Share***

One of the outliers consisted of a male and a female, of around 65 years of age who have an income from pension and other investments of around Lm5000. They spent almost Lm6000 on furniture and household equipment, which could mean that the couple has either to change some of the furniture in their home, or they bought some new equipment for the house. However, logic seems to suggest that this expense is not a normal one, i.e. for the next year the same couple will not have the same expense for this commodity.

The Kolmogorov-Smirnov test rejects the density estimates obtained.

3.5.8 ***Transport and Communication***

From a previous survey on transport carried out in Malta, an interesting result is that the average number of cars per household is more than 2. This is due to the fact that nowadays, less women are staying at home, thus these working women need to travel to go to work. Sometimes it is not feasible for their husbands to drive them around,

hence more cars are being bought, thus overpopulating the streets of this small country by cars.

The histogram is again built as before to give:

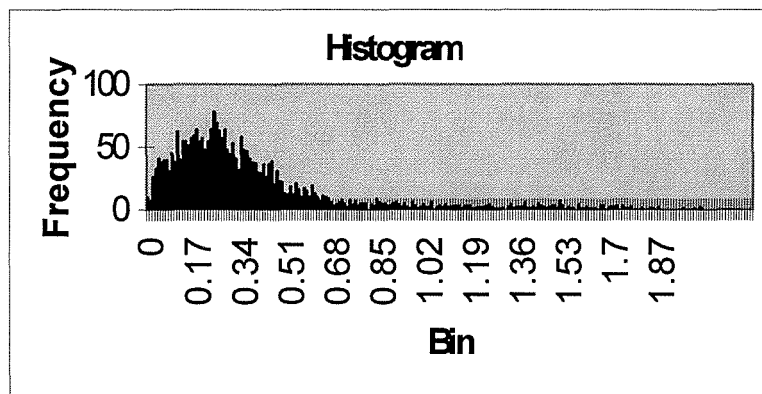


Figure 3.26 *Histogram of Transport and Communication Budget Share ($\times 10^3$)*

The shape looks like the shape of the gamma distribution, in fact the P-Plot determined that in fact, one could try to fit a gamma distribution.

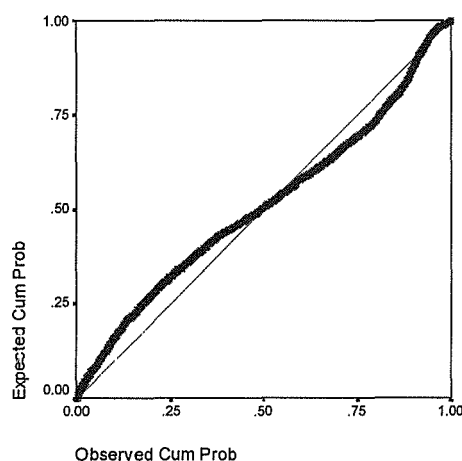


Figure 3.27 *Gamma P-Plot of Transport and Communication*

However, the gamma distribution underestimates the real data in the beginning then it overestimates them. The parameters that gave the best curve were $c = 2.5$, $b = 0.13$ and $d = 2.1$, where d is the shifting parameter. The curve obtained is as follows.

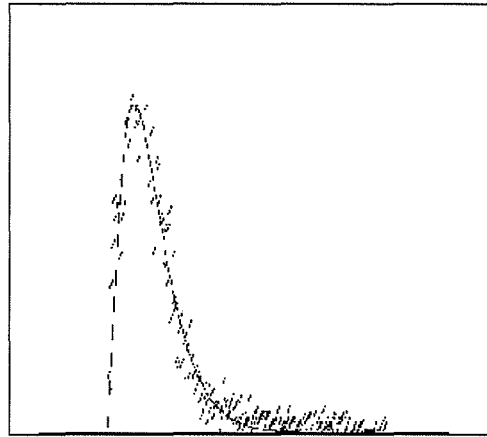


Figure 3.30 *Gamma Curve Fitted to Transport and Communication*

The equation is therefore

$$f(x) = 2.1 \left(\frac{x}{0.13} \right)^{1.5} \frac{\exp\left(\frac{-x}{0.13}\right)}{0.13\Gamma(1.5)}$$

The descriptive statistics are given by

	Minimum	Maximum	Mean		Std. Deviation	Variance
	Statistic	Statistic	Statistic	Std. Error	Statistic	Statistic
Transport & Communication	.0000	.0287	4.9019E-03	4.9353E-04	7.0491E-03	4.969E-05

Table 3.8 *Descriptive Statistics of Transport & Communication Budget Share*

All the households of the outliers spent between Lm2745 and Lm7000 in the year of the survey. A valid explanation to this is that these households bought a car during that period. That would justify the huge amount spent on this commodity. One could

verify this by looking at the expenditure of the same households in the successive year, however such data is unfortunately not available.

3.5.9 *Personal Care and Health*

Health is also becoming an issue in today's life. More people, especially young people, are investing in health insurance. Some employers are also offering beneficial conditions for their employees to enter such an insurance. The overall distribution of the health commodity is as given in the following histogram

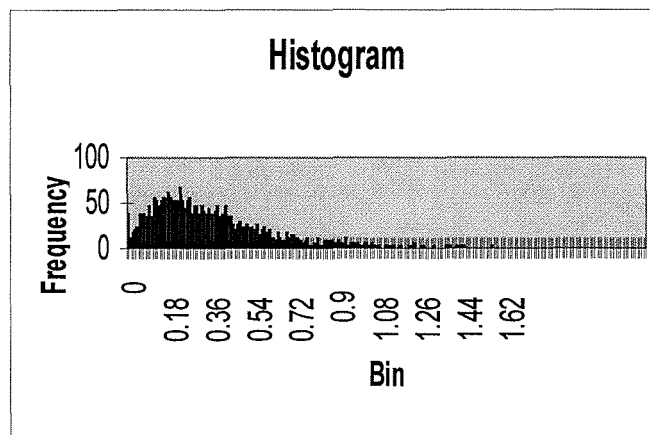


Figure 3.29 *Histogram of Personal Care and Health Budget Share($\times 10^{-3}$)*

The distribution once again looks like that of a gamma distribution. Thus before trying to fit the curve, the P-Plot was again applied to this commodity to check whether a gamma distribution could fit the data. The P-Plot confirms that a gamma distribution would fit the data perfectly.

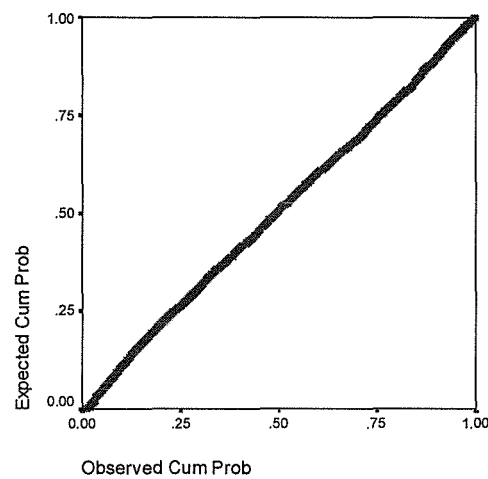


Figure 3.30 *Gamma P-Plot of Health*

The parameters fitted were $c = 2.5$, $b = 0.15$ and $d = 2$ where d is the scaling down parameter. The following diagram shows the curve fitted to the data.

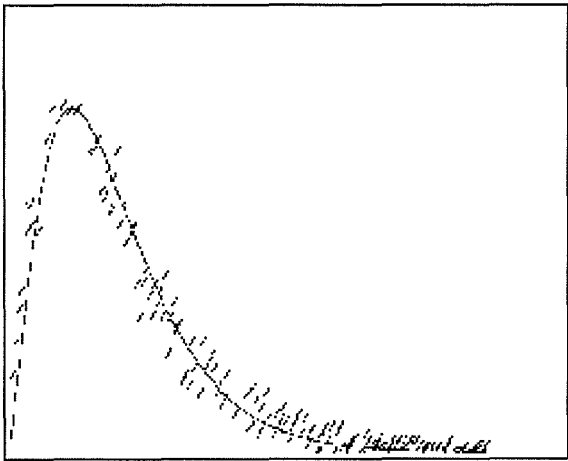


Figure 3.31 *Gamma Curve Fitted to Health*

Thus, the curve fitted is given by

$$f(x) = 2\left(\frac{x}{0.15}\right)^{1.5} \frac{\left[\exp\left(\frac{-x}{0.15}\right)\right]}{0.15\Gamma(1.5)}$$

The goodness of fit test was accepted at 0.01 level of significance.

The descriptive statistics are given by

	Minimum	Maximum	Mean		Std. Deviation	Variance
	Statistic	Statistic	Statistic	Std. Error	Statistic	Statistic
Health	.0000	.0254	5.8823E-03	5.1546E-04	6.7208E-03	4.517E-05

Table 3.9 *Descriptive Statistics of Health Budget Share*

One of the outliers is an old woman who spent around 34% of her expenditure on health, and another old woman spent around 50% on health. These women could have undergone a private operation.

3.5.10 Education and Recreation

The surveyed households were also asked how much money they spend in order to educate and recreate themselves. There are various types of hobbies a household can spend money on. An example is movie going, or night-clubs. Unfortunately, recreation is becoming an expensive hobby. Some people also spend money on education, buying books, travelling for conferences, etc. However, I would presume that this in fact is a luxury commodity and not a necessity.

Looking at the distribution of this commodity, the following diagram was obtained.

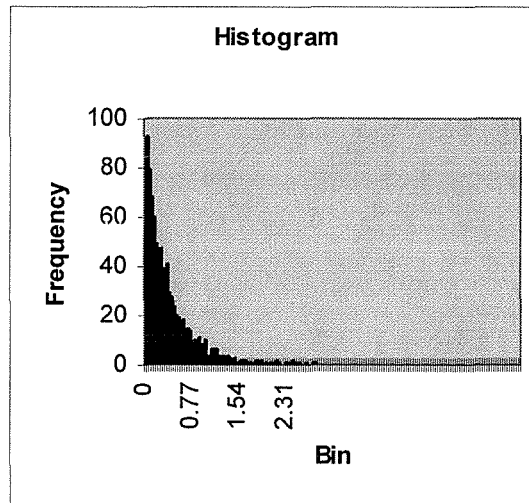


Figure 3.32 *Histogram of Education and Recreation Budget Share(*10⁻³)*

The distribution looks like an exponential distribution, and the P-Plot from the SPSS confirms this.

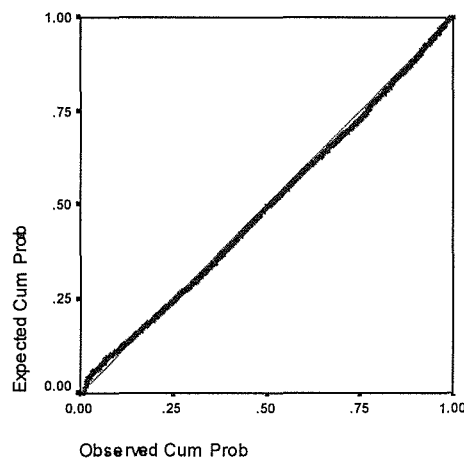


Figure 3.33 *Exponential P-Plot of Education and Recreation*

This shows that the curve in fact follows exactly an exponential distribution. Thus, data points were imported into the Qbasic program that generates an exponential distribution such that a curve will be found. Data points were multiplied by 70 for screen fitting. The fitted curve to the data points is shown in the next diagram.

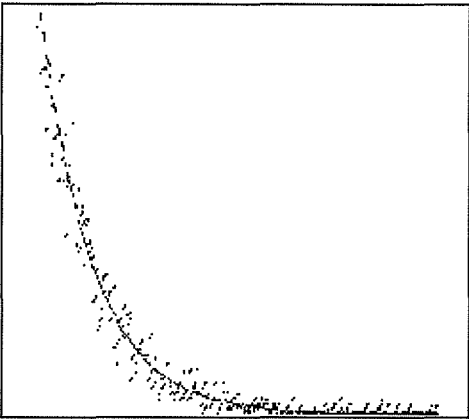


Figure 3.34 *Exponential Curve Fitted to Education and Recreation*

The parameters used are $\lambda = 2.85$ and $s = -0.1$, where s is the usual shifting parameter. Thus the curve fitted is

$$f(x - 0.1) = \frac{2.85}{70} \exp(-2.85x)$$

The goodness of fit test was accepted at 0.01 level of significance.

The descriptive statistics are given by

	Minimum	Maximum	Mean		Std. Deviation	Variance
	Statistic	Statistic	Statistic	Std. Error	Statistic	Statistic
Education & Recreation	.0000	.0343	3.5714E-03	3.7654E-04	6.3007E-03	3.970E-05

Table 3.10 *Descriptive Statistics of Education & Recreation Budget Share*

Looking at the old households, one can see that besides the outliers, all the households spend a maximum of 15% of their expenditure in education and recreation. Two of the outliers are old women living alone with an an average income of around Lm3000. As an average, they spend around 28% of their expenditure on education and recreation. If one looks now at young households, one can see that except for the outliers, all households spend less than 20% of their expenditure on education and

recreation. (Compare this with the maximum of 15% for the old households). However, three young couples spend around 30% on this commodity. One of these outliers, happens to be an outlier in the furniture and household equipment, where he also spends about 30% on that commodity. One must note that this commodity comprises also the fee for attending private schools for children. Some of these schools have a relative high fee, and nowadays, young people are more inclined to send their children to private schools, instead of the free governmental schools. If one had to study this commodity over the years, probably the average spent on this commodity would increase.

3.5.11 *Other Goods and Services*

The last commodity considered is the commodity of 'other goods and services'. This commodity consists of all the expenditures which do not fall in the previous categories (such as Pet Food, Domestic help, Dry cleaning, etc). The histogram of this commodity is given in the following diagram.

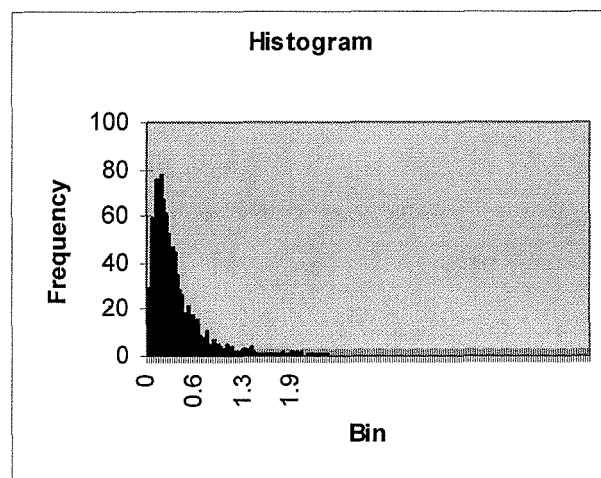


Figure 3.35 *Histogram of Other Goods and Services Budget Share($\times 10^{-3}$)*

This looks like a lognormal distribution. On applying the Plot, the following diagram was obtained.

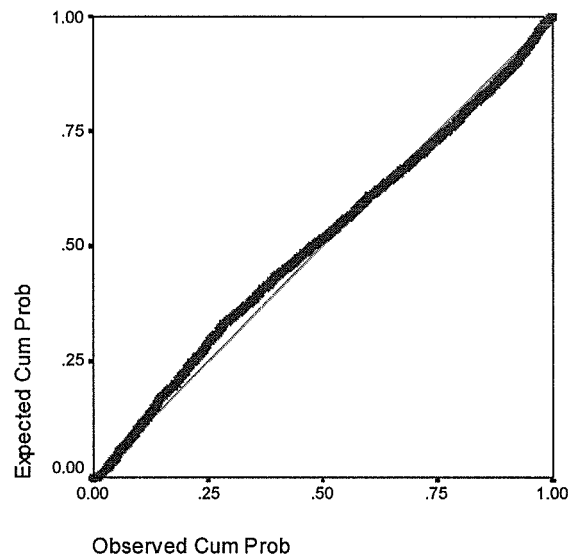


Figure 3.36 *Lognormal P-Plot of Other Goods and Services*

Thus, as in previous cases, a lognormal curve was fitted to the data points. The curve fitted had parameters $m = 0.285$, $\sigma = 0.73$, $s = 0$ and $d = 1.1$, where s is the shifting parameter in the x-axis, and d is the scaling parameter. The curve fitted to the data points is shown in the following diagram.

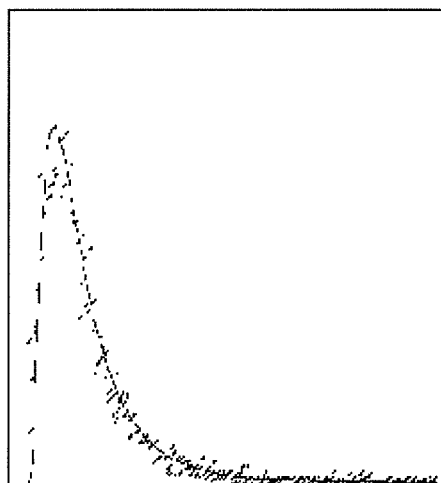


Figure 3.37 *Lognormal Curve Fitted to Other Goods and Services*

The equation is given by

$$f(x) = \frac{1.1}{x(0.73)(2\pi)^{1/2}} \exp \left\{ \frac{-\left[\log\left(x/(0.285)\right)\right]^2}{2(0.73)^2} \right\}$$

The descriptive statistics are given by

	Minimum	Maximum	Mean		Std. Deviation	Variance
	Statistic	Statistic	Statistic	Std. Error	Statistic	Statistic
Other Goods & Services	.0000	.0288	3.9840E-03	4.2514E-04	6.7354E-03	4.537E-05

Table 3.11 ***Descriptive Statistics of Other Goods & Services Budget Share***

One of the outliers is a male of 65 years who spends almost nothing on the previous commodities, however he spends more than 90% on this commodity. An explanation to this could be that the person didn't fully understand the questionnaire, and inserted all his entries in this last commodity. There is a considerable difference in the amount spent by old households, and by the amount spent by young households. All the old households spend less than 51% on this commodity. However, the young households spend less than 35% on this commodity, except for one household who spend around 55% on this commodity.

3.6 ***Improvements***

The equations obtained in the previous section could be used to simulate Maltese households, such that the previously mentioned errors concerning the data could be reduces.

One must also keep in mind that such expenditures are for local data, and therefore, the curves estimated would not apply for foreign countries. Besides, this study had to be restricted to only one year (i.e. one survey), since such surveys are not carried out

every year as in foreign countries. In Malta, such a survey is carried out once every 10 years. In my opinion, one should look at the possibility of carrying out the survey every year, thus one could look into the local time - invariance of the commodities. Other data was requested but was unavailable in Malta, such as the local alcohol and tobacco consumption. One would hope that in the future, such surveys would be carried out.

References

REFERENCES

1. **Barnett V.**, "*Interpreting Multivariate Data*", 1980, (John Wiley & Sons).
2. **Blundell R.**, "*Consumer Behaviour: Theory and Empirical Evidence – A Survey*", 1988, (The Economic Journal, No98, p16-65)
3. **Bratley P., Fox B.L. and Schrage L.E.**, "*A Guide to Simulation*", 1983, (Springer).
4. **Central Office of Statistics**, "*Household Budgetary Survey Questionnaire for 1994*", 1994.
5. **Chen K.Z.**, "*The Symmetric Problem in the Linear Almost Ideal Demand System*", 1998, (Economics Letters, No. 59, p309-315).
6. **Chen S. Xi.**, "*A Kernel Estimate for the Density of a Biological Population by Using Line Transect Sampling*", 1996, (Applied Statistics, Vol. 45, No. 2, p135-150).
7. **Christensen L.R., Jorgenson D.W. and Lau L.J.**, "*Transcendental Logarithmic Production Frontiers*", 1973, (Review Of Economics Statistics, No. 55, p23-45).
8. **Christensen L.R., Jorgenson D.W. and Lau L.J.**, "*Transcendental Logarithmic Utility Functions*", 1975, (The American Economic Review, p367-383).
9. **Deaton A. and Muellbauer J.**, "*Economics and Consumer Behavior*", 1980, (Cambridge University Press).
10. **Deaton A. and Muellbauer J.**, "*An Almost Ideal Demand System*", 1980, (The American Economic Review, Vol. 70, No. 3, p312-326).
11. **Deaton A.**, "*Understanding Consumption*", 1992, (Clarendon Press Oxford).
12. **Debreu G.**, "*Mathematical Economics*", 1986, (Cambridge University Press).
13. **Eggleston H.G.**, "*Convexity*", 1958, (Cambridge University Press).
14. **Engel J. and Kneip A.**, "*Recent Approaches to Estimating Income Distributions, Engel Curves and Related Functions*", 1994, (Discussion Paper No.A-442).
15. **Freund J.E.**, "*Mathematical Statistics*", 1992, (Prentice-Hall International).
16. **Gasser T., Muller H.G. and Mammitzsch V.**, "*Kernels for Non-parametric Curve Estimation*", 1985, (Royal Statistical Society, Vol. 47, No. 2, p238-252).

17. **Greene W.H.**, "*Econometric Analysis*", 1993, (Prentice-Hall International).
18. **Hardle W., Marron J.S. and Wand M.P.**, "*Bandwidth Choice for density Derivatives*", 1990, (Royal Statistical Society, Vol. 52, No. 1, p223-232).
19. **Hastings N.A.J. and Peacock J.B.**, "*A Handbook For Students and Practitioners*", 1975, (Butterworth & Co.(Publishers) Ltd.).
20. **Hildenbrand W.**, "*On the Empirical Evidence of microeconomic Demand Theory*", 1993, (Discussion Paper No. A-413).
21. **Hildenbrand W.**, "*How Relevant are Specifications of Behavioral Relations on the Micro-Level for Modelling the Time Path of Population Aggregates?*", 1998, (European Economic Review, No. 42, p437-458).
22. **Hildenbrand W. and Kneip A.**, "*Family Expenditure Data, Heteroscedasticity and the Law of Demand*", 1992, (Discussion Paper No. A-390).
23. **Hildenbrand W. and Kneip A.**, "*Aggregate Consumers' expenditure and Income*", 1995.
24. **Hildenbrand W. and Kneip A.**, "*Demand Aggregation under Structural Stability*", 1997, (Discussion Paper No. A-560).
25. **Hildenbrand W., Kneip A. and Utikal K.**, "*A Non-Parametric Analysis of Distributions of Household Income and Attributes*", 1998.
26. **Jerison D. and Jerison M.**, "*Measuring Consumer Inconsistency: Real Income, Revealed Preference and the Slutsky Matrix*", 1999, (Discussion Paper A-597).
27. **Lewbel A.**, "*The Rank of Demand Systems: Theory and Non-Parametric Estimation*", 1991, (Econometrica, 59, p711-729).
28. **Lluch C., Powell A.A. and Williams R.A.**, "*Patterns in Household Demand and Saving*", 1977, (International Bank).
29. **Moschini G.**, "*The Semiflexible Almost Ideal Demand System*", 1998, (European Economic Review, No. 42, p349-364).
30. **Peitz M.**, "*Utility Maximization in models of Discrete Choice*", 1995, (Discussion Paper No. A-470).
31. **Philips L.**, "*Applied Consumption Analysis*", 1974, (North-Holland Publishing Co.).
32. **Reinhard J.**, "*The Weak Axiom of Revealed Preference and Homogeneity of Demand Functions*", 1991, (Discussion Paper No. A-345).
33. **Rimmer M.T. and Powell A.A.**, "*An Implicitly Additive Demand System*", 1996, (Applied Economics, No. 28, p1613-1622).

34. **Silverman B.W.**, "*Kernel Density Estimation Using the Fast Fourier Transform*", 1982, (Applied Statistics, Vol.31).
35. **Silverman B.W.**, "*Density Estimation for Statistics and Data Analysis*", 1992, (Chapman & Hall/CRC).
36. **Scott D.W.**, "*Multivariate Density Estimation: Theory, Practice, and Visualization*", 1992, (John Wiley & Sons).
37. **Stone R.**, "*Linear Expenditure Systems and Demand Analysis: An Application to the Pattern of British Demand*", 1994, (The Economic Journal, p511-527).
38. **Thompson J.R. and Tapia R.A.**, "*Non-Parametric Function Estimation, Modeling, and Simulation*", 1990, (SIAM).
39. **Van Der Gaag J.**, "*Private Household Consumption in China*", 1985, (World Bank Staff Working Papers, Number 701).
40. **Xuezeng L., Shengming Y. and Juhuang H.**, "*The Structure of China's Domestic Consumption*", 1985, (World Bank Staff Working Papers, Number 755).

Software Used for Data Analysis

SOFTWARE USED FOR DATA ANALYSIS

1. **Borland International Inc**, "*Turbo Pascal Version 7*", 1992.
2. **Microsoft Corporation**, "*Microsoft Excel 97*", 1996.
3. **Microsoft Corporation**, "*Ms Dos Quick Basic*", 1991.
4. **SPSS Inc**, "*SPSS for Windows Version 9.01*", 1999.
5. **SPSS Inc**, "*Table Curve 2D for Windows Version 4.07*", 1996.
6. **SPSS Inc. SYSTAT Products**, "*SYSTAT Version 8*", 1998.

Appendix A

Appendix A

SOURCE CODE OF PROGRAMS WRITTEN FOR ESTIMATION

A.1 Program for Naïve Estimation written using Turbo Pascal Version 7

```
Program naive_estimator;
uses Crt;
var
  Datafile :text;
  outputfile : text;
  x : longint;
  Xi : real;
  h : real;
  value : real;
  limit : real;
  flag1 . boolean;
  flag2 : boolean;
  rewrite1 : boolean;

begin
  flag1 := false;
  flag2 := false;
  clrscr;
  write('Enter bin width');
  readln(h);
  assign(datafile,'c:\thesis\data\exp.txt');
  assign(outputfile,'c:\thesis\data\output1.txt');
  reset( datafile);
  rewrite(outputfile);
  rewrite1 := true;
  for x := 0 to 50310 do
    begin
      flag1 := false;
      flag2 := false;
      writeln(x);
      writeln(outputfile, 'x is' ,x);
      repeat
        readln(datafile,Xi);
        value := (x - Xi)/h;
        if ((value >-1) and (value < 1)) then
          begin
            writeln(outputfile,'0.5');
            flag1 := true;
          end
        else
```

```
begin
  writeln(outputfile,'0');
  if flag1 = true then
    flag2 := true;
  end;
  until (eof(datafile) or ((flag1 = true) and (flag2 = true)));
  writeln(outputfile);
  writeln(outputfile);
  x := x + trunc(h);
  if (x > 17000) and (rewrite1) then
    begin
      close(outputfile);
      assign(outputfile,'c:\thesis\data\output2.txt');
      rewrite(outputfile);
      rewrite1 := false;
    end;
  if x > 50310 then
    begin
      x:= 50310;
    end;
  reset(datafile);
end;
close(datafile);
close(outputfile);

end. {main program}
```

A.2 Curve Fitting Program for Total Expenditure Written Using Quick Basic

REM This generates a log normal function for the total expenditure

Function fitted to total expenditure

$$f(x) = \frac{1}{x(0.53)(2\pi)^{1/2}} \exp \left\{ \frac{-[\log(x/(0.34))]^2}{2(0.53)^2} \right\}$$

SCREEN 9

REM Reads data points in file a:\exp.txt

OPEN "a:\exp.txt" FOR INPUT AS #1

WINDOW (0, 0)-(2.5, 3)

PAINT (0, 3), 20

x1 = 0

pi = 22 / 7

COLOR 0

DO

INPUT #1, y1

x1 = x1 + .01

FOR a = -.01 TO .01 STEP .000025

PSET (a + x1, a + (y1 * 90))

NEXT a

LOOP UNTIL EOF(1)

COLOR 1

REM Best fitted data points

m = .34

sig = .53

FOR x = .01 TO 2 STEP .000255

REM Calculates the lognormal function

y1 = -1 * ((LOG(x / m)) ^ 2)

y2 = 2 * sig * sig

y3 = y1 / y2

y4 = EXP(y3)

```
y5 = 1 / (x * sig)
y6 = 1 / ((2 * pi) ^ .5)
y7 = y5 * y6 * y4
```

```
PSET (x, (y7))
NEXT x
```

A.3 *Curve Fitting Program for Food Budget Share Written Using Quick Basic*

REM This generates a Weibull function for the food budget share

Function fitted to food budget share

$$f(x) = 1.5 \left(\frac{2.7x^{1.7}}{0.68^{2.7}} \right) \exp \left[- \left(\frac{x}{0.68} \right)^{2.7} \right] - 0.09$$

SCREEN 9

REM Reads data points from file a:\food.txt

OPEN "a:\food.txt" FOR INPUT AS #1

WINDOW (0, 0)-(2, 1.5)

PAINT (0, 1.5), 20

COLOR 0

x1 = 0

DO

INPUT #1, y1

x1 = x1 + .01

FOR a = -.01 TO .01 STEP .000025

PSET (a + x1, a + (y1 * 70))

NEXT a

LOOP UNTIL EOF(1)

COLOR 1

REM Best fitted data points

b = 0.68

c = 2.7

s = 0.09

d = 1.5

FOR x = .01 TO 2 STEP .000255

REM Calculates the Weibull function

y2 = -1 * ((x / b) ^ c)

y3 = EXP(y2)

y4 = x ^ (c - 1)

y5 = (c * y4) / (b ^ c)

$y6 = y5 * y3$

PSET (s + x, (y6 / d))
NEXT x

A.4 *Curve Fitting Program for Beverages and Tobacco Budget Share Written Using Quick Basic*

REM This generates a Gamma function for the Beverage and Tobacco budget share

Function fitted to beverages and tobacco budget share

$$f(x) = 2 \left(\frac{x}{0.28} \right)^{1.5} \frac{\exp\left(\frac{-x}{0.28}\right)}{0.28 \Gamma(2.5)}$$

SCREEN 9

REM Reads data points from file a:\bevto.txt

OPEN "a:\bevto.txt" FOR INPUT AS #1

WINDOW (0, 0)-(4, 1)

PAINT (0, 1), 20

COLOR 0

x1 = 0

pi = 22 / 7

DO

INPUT #1, y

x1 = x1 + .01

FOR a = -.01 TO .01 STEP .00025

PSET (a + x1, a + (y * 40))

NEXT a

LOOP UNTIL EOF(1)

COLOR 1

REM Best fitted data points

c = 2.5

b = 0.28

d = 2

gamma = 1

p = c

FOR x = .005 TO 5 STEP .0005

st = c

DO

st = st - 1

gamma = gamma * st

```
IF st < 1 THEN st = 0  
LOOP UNTIL st = 1
```

```
y1 = x / b  
y2 = (y1) ^ ((p + .5) - 1)  
y3 = EXP(-1 * y1)  
y4 = y2 * y3  
y5 = b * (gamma * (pi ^ .5))  
y6 = y4 / y5
```

```
PSET (x, y6 / d)  
NEXT x
```

A.5 Curve Fitting Program for Clothing and Footwear Budget Share Written Using Quick Basic

REM This generates an Exponential function for the clothing and footwear budget share

Function fitted to clothing and footwear budget share

$$f(x) = 2.2 \exp(-2.2x) + 0.22$$

SCREEN 9

REM Reads data points from file a:\clothes.txt

OPEN "a:\clothes.txt" FOR INPUT AS #1

WINDOW (0, 0)-(2, 1.5)

PAINT (0, 1.5), 20

COLOR 0

x1 = 0

DO

INPUT #1, y1

x1 = x1 + .01

FOR a = -.01 TO .01 STEP .000025

PSET (a + x1, a + (y1 * 70))

NEXT a

LOOP UNTIL EOF(1)

COLOR 1

REM Best fitted data points

lambda = 2.2

s = -0.22

FOR x = 0 TO 2.2 STEP .00005

y = lambda * (EXP(-lambda * x))

PSET (s + x, (y))

NEXT x

INPUT s

A.6 Curve Fitting Program for Housing Budget Share Written Using Quick Basic

REM This generates an Exponential function for the housing budget share

Function fitted to the housing budget share

$$f(x) = 3.2 \exp(-3.2x) + 0.15$$

SCREEN 9

REM Reads data points from file a:\housing.txt

OPEN "a:\housing.txt" FOR INPUT AS #1

WINDOW (0, 0)-(3, 1.2)

PAINT (0, 1.2), 20

COLOR 0

x1 = 0

DO

INPUT #1, y1

x1 = x1 + .01

FOR a = -.01 TO .01 STEP .000025

PSET (a + x1, a + (y1 * 70))

NEXT a

LOOP UNTIL EOF(1)

COLOR 1

REM Best fitted data points

lambda = 3.2

s = -0.15

FOR x = 0 TO 3 STEP .00025

y = 1 * (EXP(-1 * x))

PSET (s + x, (y))

NEXT x

A.7 *Curve Fitting Program for Fuel and Power Budget Share Written Using Quick Basic*

REM This generates a Lognormal function for the fuel and power budget share

Function fitted to fuel and power budget share

$$f(x) = \frac{1.1}{x(0.65)(2\pi)^{1/2}} \exp \left\{ \frac{-[\log(x/(0.3))]^2}{2(0.65)^2} \right\}$$

SCREEN 9

REM Reads data points from file a:\fuel.txt

OPEN "a:\fuel.txt" FOR INPUT AS #1

WINDOW (0, 0)-(2.5, 3)

PAINT (1, 3), 20

COLOR 0

x1 = 0

pi = 22 / 7

DO

INPUT #1, y1

x1 = x1 + .01

FOR a = -.01 TO .01 STEP .000025

PSET (a + x1, a + y1 * 90)

NEXT a

LOOP UNTIL EOF(1)

COLOR 1

REM Best fitted data points

m = 0.3

sig = 0.65

s = 0

d = 1.1

FOR x = .01 TO 3 STEP .00005

y1 = -1 * ((LOG(x / m)) ^ 2)

y2 = 2 * sig * sig

y3 = y1 / y2

y4 = EXP(y3)

y5 = 1 / (x * sig)

```
y6 = 1 / ((2 * pi) ^ .5)  
y7 = y5 * y6 * y4
```

```
PSET (s + x, (y7 / d))  
NEXT x
```

A.8 Curve Fitting Program for Furniture and Household Equipment Budget Share Written Using Quick Basic

REM This generates a Gamma function for the furniture and household equipment budget share

Function fitted to furniture and household equipment budget share

$$f(x) = 0.07 + \left(\frac{x}{1.5}\right)^{0.5} \frac{\exp\left(\frac{-x}{1.5}\right)}{1.5\Gamma(0.5)}$$

SCREEN 9

REM Reads data points from file a:\furnhou.txt

OPEN "a:\furnhou.txt" FOR INPUT AS #1

WINDOW (0, 0)-(3, 1)

PAINT (0, 1), 20

COLOR 0

x1 = 0

pi = 22 / 7

DO

INPUT #1, y

x1 = x1 + .01

FOR a = -.01 TO .01 STEP .000025

PSET (a + x1, a + (y * 90))

NEXT a

LOOP UNTIL EOF(1)

COLOR 1

REM Best fitted data points

c = 1.5

b = 1.5

d = -0.07

p = c

FOR x = .005 TO 7 STEP .00005

pi = 22 / 7

y1 = x / b

y2 = (y1) ^ (-.5)

y3 = EXP(-1 * y1)

y4 = y2 * y3

y5 = b * (pi ^ .5)

```
y6 = y4 / y5  
PSET (x+d, y6)  
NEXT x
```

A.9 Curve Fitting Program for Transport and Communication Budget Share Written Using Quick Basic

REM This generates a Gamma function for the transport and communication budget share

Function fitted to transport and communication budget share

$$f(x) = 2.1 \left(\frac{x}{0.13} \right)^{1.5} \frac{\exp\left(\frac{-x}{0.13}\right)}{0.13\Gamma(1.5)}$$

SCREEN 9

REM Reads data points from file a:\trancom.txt

OPEN "a:\trancom.txt" FOR INPUT AS #1

WINDOW (-.5, 0)-(2.5, 1)

PAINT (0, 1), 20

COLOR 0

x1 = 0

pi = 22 / 7

DO

INPUT #1, y

x1 = x1 + .01

FOR a = -.01 TO .01 STEP .000025

PSET (a + x1, a + (y * 40))

NEXT a

LOOP UNTIL EOF(1)

COLOR 1

REM Best fitted data points

c = 2.5

b = 0.13

d = 2.1

p = c

FOR x = .005 TO 5 STEP .00005

gamma = 1

st = c

pi = 22 / 7

DO

```
st = st - 1  
gamma = gamma * st  
IF st < 1 THEN st = 0  
LOOP UNTIL st = 1
```

```
y1 = x / b  
y2 = (y1) ^ ((p + .5) - 1)  
y3 = EXP(-1 * y1)  
y4 = y2 * y3  
y5 = b * (gamma * (pi ^ .5))  
y6 = y4 / y5
```

```
PSET (x, y6 / d)  
NEXT x
```

A.10 Curve Fitting Program for Personal Care and Health Budget Share Written Using Quick Basic

REM This generates a Gamma function for the Personal Care and Health budget share

Function fitted to personal care and health budget share

$$f(x) = 2 \left(\frac{x}{0.15} \right)^{1.5} \frac{\exp\left(\frac{-x}{0.15}\right)}{0.15\Gamma(1.5)}$$

SCREEN 9

REM Reads data points from file a:\health.txt

OPEN "a:\health.txt" FOR INPUT AS #1

WINDOW (0, 0)-(2, 1)

PAINT (0, 1), 20

COLOR 0

x1 = 0

pi = 22 / 7

DO

INPUT #1, y

x1 = x1 + .01

FOR a = -.01 TO .01 STEP .000025

PSET (a + x1, a + (y * 40))

NEXT a

LOOP UNTIL EOF(1)

COLOR 1

REM Best fitted data points

c = 2.5

b = 0.15

d = 2

p = c

FOR x = .005 TO 1.7 STEP .00005

gamma = 1

st = c

pi = 22 / 7

DO

st = st - 1

gamma = gamma * st

IF st < 1 THEN st = 0

LOOP UNTIL st = 1

y1 = x / b

y2 = (y1) ^ ((p + .5) - 1)

y3 = EXP(-1 * y1)

y4 = y2 * y3

y5 = b * (gamma * (pi ^ .5))

y6 = y4 / y5

PSET (x, y6 / d)

NEXT x

A.11 Curve Fitting Program for Education and Recreation Budget Share Written Using Quick Basic

REM This generates an Exponential function for the Education and Recreation Budget Share

Function fitted to education and recreation budget share

$$f(x) = 2.85 \exp(-2.85x) + 0.1$$

SCREEN 9

REM Reads data points from file a:\edurec.txt

OPEN "a:\edurec.txt" FOR INPUT AS #1

WINDOW (0, 0)-(3, 1.5)

PAINT (0, 1.5), 20

COLOR 0

x1 = 0

DO

INPUT #1, y1

x1 = x1 + .01

FOR a = -.01 TO .01 STEP .000025

PSET (a + x1, a + (y1 * 70))

NEXT a

LOOP UNTIL EOF(1)

COLOR 1

REM Best fitted data points

lambda = 2.85

s = -0.1

FOR x = 0 TO 2.8 STEP .0005

y = lambda * (EXP(-lambda * x))

PSET (s + x, (y))

NEXT x

REM NEXT col

INPUT s

A.12 Curve Fitting Program for Other Goods and Services Budget Share Written Using Quick Basic

REM This generates a Lognormal function for the Other Goods and Services Budget Share

Function fitted to other goods and services budget share

$$f(x) = \frac{1.1}{x(0.73)(2\pi)^{1/2}} \exp \left\{ \frac{-[\log(x/(0.285))]^2}{2(0.73)^2} \right\}$$

SCREEN 9

REM Reads data points from file a:\other.txt

OPEN "a:\other.txt" FOR INPUT AS #1

WINDOW (-.1, 0)-(2.9, 3)

PAINT (1, 3), 20

COLOR 0

x1 = 0

pi = 22 / 7

DO

INPUT #1, y1

x1 = x1 + .01

FOR a = -.01 TO .01 STEP .000025

PSET (a + x1, a + y1 * 90)

NEXT a

LOOP UNTIL EOF(1)

COLOR 1

REM Best fitted data points

m = 0.285

sig = 0.73

s = 0

d = 1.1

FOR x = .01 TO 3 STEP .00005

y1 = -1 * ((LOG(x / m)) ^ 2)

y2 = 2 * sig * sig

y3 = y1 / y2

```
y4 = EXP(y3)
y5 = 1 / (x * sig)
y6 = 1 / ((2 * pi) ^ .5)
y7 = y5 * y6 * y4
```

```
PSET (s + x, (y7 / d))
NEXT x
```