

UNIVERSITÉ CATHOLIQUE DE LOUVAIN
Institut de Statistique, Biostatistique et Sciences Actuarielles

Statistical inference in efficiency analysis with applications

Thèse présentée en vue de l'obtention du grade de
Docteur en Sciences (orientation statistique) par :
Rachida El Mehdi

Membres du Jury :

Prof. Anouar El Gouch (*Secrétaire*)

Prof. Bachir Elkihel (*Co-Promoteur*)

Prof. Christian Hafner (*Promoteur*)

Prof. Léopold Simar

Prof. Ingrid Van Keilegom (*Président du jury*)

Prof. Philippe Vanden Eeckaut

Louvain-la-Neuve

Septembre 2015

To my mother

Acknowledgements

First of all, I owe my deepest gratitude to my supervisor Professor Christian M. Hafner who agreed to lead this thesis. I thank him for his motivation, availability, patience and for the knowledge acquired since the preparation of my Post-Graduate Diploma (DEA). His pedagogy, meticulous character and his orientations had guided me a lot in the realization and the writing of my works. He always solved kindly any problem at the research level or at my doctoral process level. I thank him infinitely for the repetitive meetings and the precious time which he had reserved to me. The collaboration with him was, is and will be an immense pleasure for me.

My thanks to the co-supervisor of this thesis Professor Bachir Elkihel for his assistance, helps and remarks. He was always available to propose me data and has never stopped to propose me conferences and colloques in order to present my works. I hope that we can realize, in a near future, an efficiency analysis on data of the Moroccan industry.

Sincere thanks to Professor Léopold Simar for his great help, constructive comments and suggestions. I thank him to agree to be a member of the accompaniment committee of my thesis and for the long meetings and discussions that he had granted me alone or in the presence of Professor Hafner in spite of his very charged schedule. Thank you for paying attention to my presentations and directing me in the followed procedures. Certainly, it is an honor for me to meet him.

I would like to thank Professor Anouar EL Ghouch for his supports, encouragements and advices. His office was always open for me, with or without appointment, to answer my technical or computer questions. Thank you also for agreeing to be a member of the accompaniment committee and for having repeatedly spent time in the computer room to revise and correct my programs. He finds here all my respect and my great recognition.

Besides the accompaniment committee, I am grateful to the Professors Ingrid Van Keilegom and Philippe Vanden Eeckaut to accept to be part of the doctoral jury. It is also the opportunity to thank all the professors of the Institute of Statistics, Biostatistics and Actuarial Sciences with whom I attended a course, a short course or a seminar within the framework of the doctoral training. Their

courses and works were a source of inspiration and very precise references.

I am indebted to the Belgian Technical Cooperation (BTC or CTB) for the funds allocated for this research. It had allowed me to prepare my PhD thesis in good conditions. Its staff in Belgium and particularly Christine Leroy, Aurélie Vandecruys and Céléstin Misigaro and the staff of its representation in Morocco and principally Amal Hadaaj had all friendly welcomed me and had facilitated me all the procedures. It is also an opportunity to appreciate and to thank enormously the Catholic University of Louvain (UCL) for supporting the end of my doctoral program and to offer me an appropriate working environment.

I am indebted too to the Mohammed First University of Oujda and mainly its National School of Applied Sciences for the free time assigned to the realization of this work. I greet its spirit of cooperation and its opening on the rest of the world. Thanks to every person who helped me in my steps and initiatives.

Many thanks also go to my colleagues. In Morocco, I thank them all without exception for offering me an adequate frame for the work and the research. Their friendship and their encouragements accompanied me throughout my road. In Belgium, I thank pleasant administrative staff for its assistance and its understanding; the respectful IT staff, without forgetting M. Jean Luc Marrion, who had insured me a very favorable framework for the execution of my codes being spot or remote and who resolved my local problem of remote connection to the computers of the researchers room; all the assistants and the researchers who had the kindness to answer my questions and to not interrupt the execution of my codes when I used a distant computer. I was welcomed always warmly in this family atmosphere which reigns at the institute.

Thanks to the organizers of the PAI, the MAMERN'11 conference of Saïdia, the CIGIMS2012 congress of Fez, the CEAFE12 conference of Cairo and the summer school of Pisa in 2007 for the acquired knowledge and finally thanks to the organizers of any scientific meeting which had a relation with my research and who accepted the presentation of my works but I was not able to attend the event for one reason or another.

Last but not the least, I would like to thank my family, my friends and those who made this thesis possible. So, thank you for all.

Abstract

This thesis is a contribution to frontier analysis and its application to developing areas in Morocco. It focuses on both nonparametric and parametric frontier methods and mainly proposes inference methods for efficiency under weak conditions. In a stochastic frontier analysis, a possible dependence between the two error terms is considered for both cross-sectional and panel data. In order to validate our models, several tests are performed and analyses of the sensitivity to outliers are established.

Confidence intervals are constructed using the smoothed bootstrap for the nonparametric approach and two extended parametric percentile bootstrap algorithms for the parametric one. Both are illustrated using simulated and real data. For panel data, a variety of alternative deterministic time effect models are proposed, allowing e.g. for periodicity and disappearing inefficiency over time. Efficiency scores generally suggest the inefficiency of the handled sectors, but results should be interpreted carefully given the unavailability of certain factors.

Contents

1	Introduction	1
1.1	The frontier concept	3
1.1.1	The deterministic and the stochastic nonparametric approaches	10
1.1.2	The deterministic and the stochastic parametric approaches	21
1.2	What is local government in Morocco?	22
1.3	The drinking water sector in Morocco	22
1.4	The data motivation	23
1.5	Comparison between nonparametric and parametric approaches	24
2	Local government efficiency: The case of Moroccan municipalities	29
2.1	Introduction	29
2.2	The DEA program	31
2.3	Inference using the bootstrap	33
2.3.1	Data Generating Process	34
2.3.2	Bootstrap correcting bias for DEA efficiency scores	34
2.3.3	Confidence intervals for DEA efficiency scores	35
2.3.4	Testing returns to scale	36
2.4	The free disposal hull approach	37
2.5	Illustrative application to local government efficiency in Morocco	39
2.5.1	The data	40
2.5.2	Interpretation of the results	42
2.6	Conclusion	50

3 Inference in stochastic frontier analysis with dependent error terms	53
3.1 Introduction	53
3.2 Parametric stochastic frontier models	55
3.2.1 Classical SFA with independence	56
3.2.2 SFA with dependent error components	57
3.2.3 Bootstrap confidence intervals for technical efficiencies	59
3.3 An illustrative analysis of the efficiency of Moroccan municipalities	62
3.4 Conclusion	72
4 Statistical inference for panel stochastic frontier analysis with an illustration of Moroccan drinking water performance	75
4.1 Introduction	75
4.2 Efficiency measures for panel data	77
4.2.1 The panel data production frontier model	78
4.2.2 Time-varying efficiency	79
4.2.3 Model estimation	81
4.3 Inferences on the Technical Efficiency measure	85
4.4 Illustrative empirical analysis on Moroccan drinking water area	87
4.5 Conclusion	95
5 Conclusions	97
Appendices	104
A Some definitions and proofs for the DEA approach	105
B Some definitions and properties of copula functions	109
B.1 Gaussian copula	110
B.2 FGM copula	111
B.3 Ali-Mikhail-Haq (AMH) copula	112
B.4 Clayton copula	112
B.5 Frank copula	113

C MLE with independent error terms and time-varying technical efficiency	115
C.1 The time-constant efficiency	115
C.2 Maximum Likelihood Estimation for the model with independent errors terms and time-varying technical efficiency	116
D R tools used	119
D.1 R packages	119
D.2 Estimated computation time	121
E Positive skewness in the panel stochastic frontier analysis	123
E.1 The extended exponential distribution	123
E.1.1 The density of the error	124
E.1.2 Technical Efficiency expression	125
E.2 The extended half-normal distribution	126
Bibliography	127

List of Figures

1.1	A DEA and a FDH nonparametric frontier for a set of points . .	7
1.2	The DEA and SFA frontiers comparison	7
1.3	Technical and allocative efficiencies	9
2.1	Financial autonomy versus operating receipts and DEA frontier	41
2.2	Financial autonomy versus operating receipts and DEA frontier (Zoom)	42
2.3	Graph of bias corrected estimators and their confidence interval (Districts ranked according to increasing population size)	49
3.1	Boxplot of TE for ten models with HN and TN($\mu = -1$)	70
3.2	$g_{\theta}(\varepsilon)$ distribution for the N-HN-Clay Copula model	71

List of Tables

2.1	Efficiency scores of FDH, DEA with VRS and their confidence intervals	44
2.1	Efficiency scores of FDH, DEA with VRS and their confidence intervals	45
2.1	Efficiency scores of FDH, DEA with VRS and their confidence intervals	46
2.1	Efficiency scores of FDH, DEA with VRS and their confidence intervals	47
3.1	Estimated coverages of $E[e^{-U} \theta, x_0 = 60, y_0 = y_{(q)}]$ by bootstrap confidence intervals.	62
3.2	Technical Efficiency and Log-Likelihood values in the independence case	64
3.3	Maximum likelihood estimates in the independence case	65
3.4	Technical Efficiency and Log-Likelihood values with dependence	67
3.5	ML estimator of ϑ for the HN-Clay model	69
3.6	TE confidence intervals for the HN-Clay model	72
4.1	The AIC values of the estimated models	90
4.2	The model estimation with correlated error terms	90
4.3	Technical Efficiency scores for the 2001-2007 period and their confidence intervals	92
4.3	Technical Efficiency scores for the 2001-2007 period and their confidence intervals	93

4.4	Provinces number in each region according to the mean of TE estimates	94
-----	--	----

Chapter 1

Introduction

Measuring efficiency is useful for any Decision Making Unit (DMU) to be situated with regard to the other institutions which have a purpose to produce one or several outputs from a certain number of inputs. This comparison allows the DMU to judge the quality of its management and to improve it, when it is possible, by adopting the policies of the most efficient entities. The efficiency analysis can be done, hence, in all domains related to output-input relationship.

Hence, measuring efficiency is essential because first it allows to determine if the DMU succeeds or fails to reach the efficiency and to compare the DMU's performances; second, if inefficiency is observed, the causes of this inefficiency can be identified and therefore eliminated. This helps the inefficient DMU to develop an optimal policy which is superior to the current one.

Since the 1950s, several studies have been made worldwide on the efficiency measure. The pioneering works in this sense were developed by Koopmans (1951), Debreu (1951) and Farrell (1957). The peculiarity of this research is to handle Moroccan data and to model the dependence between the two components of the error term in the efficiency analysis and mainly to make the inference on these measures given that they are estimated values.

Indeed, many research areas have been subject to frontier analysis, we refer

notably to the recent works on public spending as the public services in Perelman (1996); the healthcare in the world in Greene (2005) and in the Lovell (2006) works; the management quality and the productivity of the hospitals in Finland, England and the Nordic countries in Linna (1998), Jacobs (2001) and Linna et al. (2010); the education in developing countries, Australian universities and English and Welsh universities in Greene (2005), Abbott and Doucouliagos (2003, 2009) and Stevens (2005) respectively; and the local government performance in Afonso and Fernandes (2008), Moore et al. (2005) and Borger and Kerstens (1996) for respectively Portuguese municipalities, largest cities in the United States and Belgian municipalities¹. We also refer to the energy application as in Haney and Pollitt (2009) about forty countries in the world and the electricity field as the Hirschhausen, Cullmann and Kappeler (2006) work on the German data, the comparative study Hattori, Jamasb and Pollitt (2005) between UK and Japan, the Rosen, Le and Dincer (2005) paper for the Canadian city of Edmonton and also the Ambapour (2001) work on some African producers of electricity. Other examples exist as the De Witte and Marques (2008), the Tupper and Resende (2004) and the Cubbin and Tzanidakis (1998) works for the industry water in four European countries and Australia, in Brazil and in England and Wales respectively; the Erber (2006) work on the telecommunication industries in the US and major European countries; the Agahi, Zarafshani and Behjat (2008) on crop insurance of Wheat Farmers in Kermanshah Province in Iran; the Coelli (1995b) and the Thompson et al. (1990) concerning the agriculture in Australia and in the American Kansas city; the Coelli and Perelman (1999) and the Mbangala and Perelman (1997) studies on the European and the Sub-Saharan Africa railways.

This is merely a selective overview of some researches published in frontier analysis. Hence it is not an exhaustive list of all recent papers in the various domains, but an overview which however leads us to deduct that there is a lack of efficiency research in the underdeveloped and developing countries in comparison with the developed ones. For instance, Worthington and Dollery (2002) indicates that the use of the Data Envelopment Analysis (DEA) as a

¹The work Holzer et al. (2009) gives a literature review and analysis related to measurement of local government efficiency.

technique for measuring the efficiency of government service delivery is relatively well established in Australia and several other advanced countries.

In this sense, several public or private Moroccan entities are susceptible in an efficiency analysis but it seems, to the best of our knowledge, that unfortunately no analysis has been made in this framework until now. The lack of interest of the responsible people in the number, the unavailability of the data and their confidentiality in certain cases can be the reasons which block the research at this level. Often, even if certain data exist, other data, which would allow us to deepen the analysis, are missing. So, we will try in this work to treat certain domains as illustrative examples such as the financing of local authorities and the production of drinking water in Morocco.

The choice of these sensitive sectors is not an arbitrary matter but it imposes itself. On the one hand, the financing of local authorities because I participated for a long time in the elaboration of these data which were analyzed by simple tools of descriptive statistics. The sector of the production of water imposes itself due to the interest that it carries for active persons in the environment sector and due the shortage in water noticed in Morocco and worldwide. The availability of some data in these domains suggests to use them first and foremost. Certainly, there are other data sets which do not require a lot of effort to be collected and analyzed and which could, potentially, reveal other research horizons, we hope to study some of them in future research. So, we will in the following introduce the frontier analysis and give an overview of the local government, the drinking water area in Morocco and their data motivation.

Whatever the studied sector and whatever the approach used, the performance is measured by the efficiency in the frontier analysis. It turns out so essential such that, before beginning any analysis, we will define the frontier concept. What does it mean then?

1.1 The frontier concept

The frontier notion appoints a limit function which envelops a points set. In other words, the frontier is a kind of envelope which often coincides with all points identified as representative of the best practice in the production field,

and with regard to which, the performance of every company can be compared. It delimits hence the feasible domain for all the DMUs under study. It was introduced in Farrell (1957) which defines the measurement of productive efficiency and estimates a set enveloping the cloud of data points and the resulting efficiency scores. These Farrell's efficiencies have a similarity with the coefficients of resource utilization of Debreu (1951). Later, the concept was largely developed by Charnes et al. (1978) and Banker et al. (1984) and others, proposing hence new tools in the frontier analysis.

To define the frontier two approaches are adopted. The nonparametric approach which uses mainly the Data Envelopment Analysis (DEA) method or the Free Disposal Hull (FDH) method and the parametric approach which uses particularly the Stochastic Frontier Analysis (SFA) method. Before explaining these various approaches in the frontier analysis, let us define first of all the returns to scale notion, its different cases, the Cobb-Douglas and the translog functions and the efficiency notion.

It is often assumed that the production possibility set, denoted Ψ and defined as $\Psi = \{(x, y) \mid x \text{ can produce } y\}$, is convex such as for input quantities x_1, x_2 and x and output quantities y_1, y_2 and y , if two points (x_1, y_1) and (x_2, y_2) are in Ψ , then for all $\alpha \in [0, 1]$ the point $(x, y) = \alpha(x_1, y_1) + (1 - \alpha)(x_2, y_2)$ is also in Ψ . On the other hand, the production set boundary or the technology frontier, denoted Ψ^∂ and defined by $\Psi^\partial = \{(x, y) \in \Psi \mid (\theta x, y) \notin \Psi, \forall 0 < \theta < 1, (x, \gamma y) \notin \Psi, \forall \gamma > 1\}$, is assumed to display locally either constant returns to scale (CRS), increasing returns to scale (IRS) or decreasing returns to scale (DRS). If Ψ^∂ has different regions that display IRS, CRS and DRS, respectively, then Ψ^∂ is said to display variable returns to scale (VRS). Hence, returns to scale describes what happens as the scale of production is increased. It explains the behavior of the rate of increase in output relative to the associated increase in the inputs, see Simar and Wilson (2015) for additional details.

Generally, the production function of the firm could exhibit different types of returns to scale in different ranges of output. Typically, there could be IRS at relatively low output levels, DRS at relatively high output levels, and CRS at one output level between those ranges. The IRS is the situation where the

output increases by more than the proportional change in inputs; the DRS is that where the output increases by less than the proportional change in inputs; and the CRS is that where the output increases by the same as the proportional change in inputs. This means that for a technology f which uses p inputs $x_1, \dots, x_p$ and for a constant a ,

$$IRS : \forall a > 1, \quad f(ax_1, \dots, ax_p) > af(x_1, \dots, x_p) \quad (1.1.1)$$

$$DRS : \forall a > 1, \quad f(ax_1, \dots, ax_p) < af(x_1, \dots, x_p) \quad (1.1.2)$$

$$CRS : \forall a > 0, \quad f(ax_1, \dots, ax_p) = af(x_1, \dots, x_p) \quad (1.1.3)$$

The associated formulas of these different regions of returns to scale in the frontier analysis will be provided in Section 2.2 of Chapter 2.

Let us also give an overview of the Cobb-Douglas and the translog functions which are two particular functional forms of the production function. In the input-output space, the basic form of the Cobb-Douglas frontier function when variables are in the log scale and when p variables $x_1, \dots, x_p$ are used to produce one output y and when there are n individuals under study, is as follows:

$$y_i = \beta_0 + \sum_{j=1}^p \beta_j x_{ij} + \epsilon_i, \quad (1.1.4)$$

for all $i = 1, \dots, n$ and where β_0 and $\beta_1, \dots, \beta_p$ are parameters and ϵ_i is an error term. Due to its simplicity, the Cobb-Douglas function is the most used in the estimation of the frontier function. However, this simplicity is associated with a number of restrictions on e.g. the returns to scale and the elasticity² of substitution of the factors which are constant for the Cobb-Douglas function. Hence, the model rests on some very specific hypotheses.

An alternative function is the translog one defined in Christensen, Jorgenson and Lau (1971) and which is the flexible form the most used. It imposes no restrictions upon returns to scale or substitution possibilities. Even if the flexible form presents the advantage to be able to represent any technology, it has

²Elasticity is a measure of a variable's relative sensitivity to a relative change in another variable. For two variables x_j and $x_{j'}$, its simplest form can be given as $E_{x_{j'}|x_j} = (\Delta x_{j'}/x_{j'}) / (\Delta x_j/x_j)$

however some limits as it cannot fully satisfy the regularity conditions satisfied by the Cobb-Douglas function. The simplest form of the translog function in the cross-sectional case is defined as

$$y_i = \beta_0 + \sum_j \beta_j x_{ij} + \sum_j \sum_{j'} \beta_{jj'} x_{ij} x_{ij'} + \epsilon_i, \quad (1.1.5)$$

for all $i = 1, \dots, n$ and $j, j' = 1, \dots, p$. Hence, the Cobb-Douglas function is a particular case of the translog one.

Figure 1.1³ represents an example of the nonparametric frontier in the case of the DEA and the FDH for a set of points of one input and one output. It reflects three types of envelopes appointing three frontiers: the most restrictive one indicates the DEA frontier in the presence of the constant return to scale, noted DEA-CRS, the second one indicates the DEA frontier in the case of the variable return to scale denoted DEA-VRS, and finally the dashed one which is the least restrictive one indicates the FDH frontier which ignores the convexity of the envelope curve.

Indeed, the half-line which passes by the origin and the point C draws the constant frontier associated to the DEA-CRS which indicates that only the DMU represented by the point C is efficient. On the other hand, the variable frontier associated to the DEA-VRS indicates that two other situations next to that of C which are A and E are efficient also and so, any point situated on this frontier is considered as efficient. Besides, being less restrictive, the envelope describing the FDH frontier shows that only the point T, situated under the frontier, is inefficient. So, the DMU represented by B is declared efficient by FDH but inefficient by both DEA-VRS and DEA-CRS. As for the parametric approach, the determination of the frontier is based on the estimation of the log-likelihood function. It supposes in its general case that the distances with regard to the frontier are either an inefficiency or a random shock. Usually, the production function in this approach is a Cobb-Douglas or a translog function.

³Graph from the work of Jean Bourdon entitled "La mesure de l'efficacité scolaire par la méthode de l'enveloppe : test des filières alternatives de recrutement des enseignants dans le cadre du processus Education pour tous" In the 26th Days of applied Microeconomy, France (2009) [halshs-00399562 - version 1].

Figure 1.1: A DEA and a FDH nonparametric frontier for a set of points

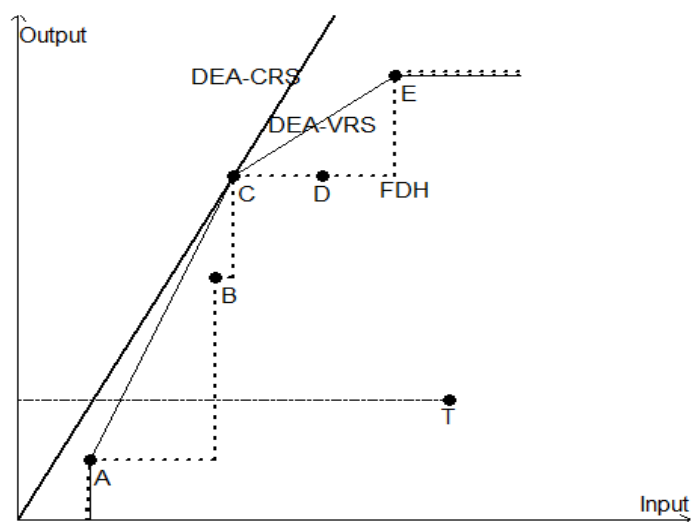
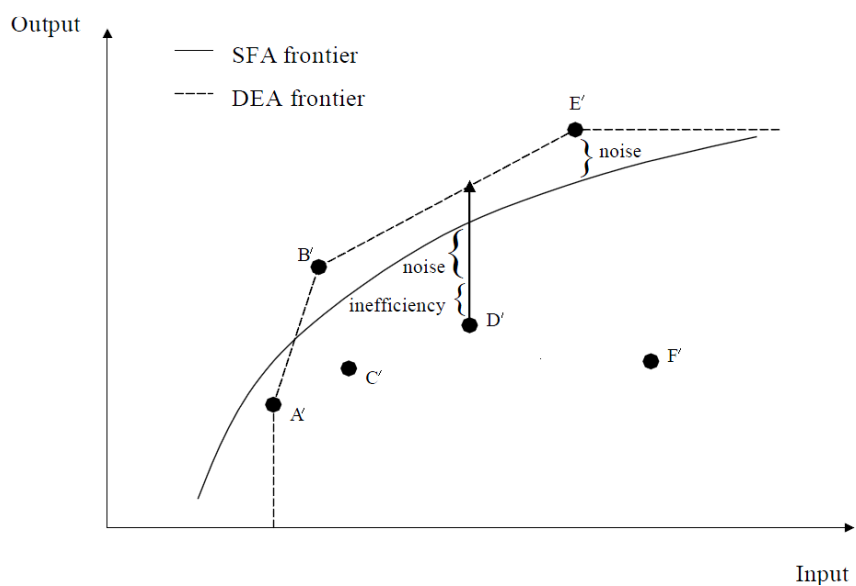


Figure 1.2: The DEA and SFA frontiers comparison



To illustrate the comparison between the DEA frontier and that of the SFA we consider Figure 1.2⁴ in which are represented the input and the output combinations for a set of DMUs. The dotted and the solid curves are respectively the DEA and the SFA frontiers and indicate the maximum output that could be produced for each level of input. The figure shows that the A' unit is DEA efficient and SFA inefficient. For the SFA inefficient unit D' , two types of deviation are noticed, the statistical noise and the inefficiency; it is also DEA inefficient but in this case all the deviation from the DEA frontier is an inefficiency. The E' point is both DEA and SFA efficient; when we assume that its inefficiency is null, the distance between the two frontiers represents a noise.

Since its introduction in 1957, Farrell's efficiency measure has been widely used in empirical researches in order to measure the efficiency of firms, countries, or other decision making units, see e.g. Färe (1984) for a detailed survey. The literature presents mainly two types of the efficiency measures which are the technical efficiency and the allocative efficiency. They will be clarified for the case of a firm which uses two inputs to produce one output under constant returns to scale.

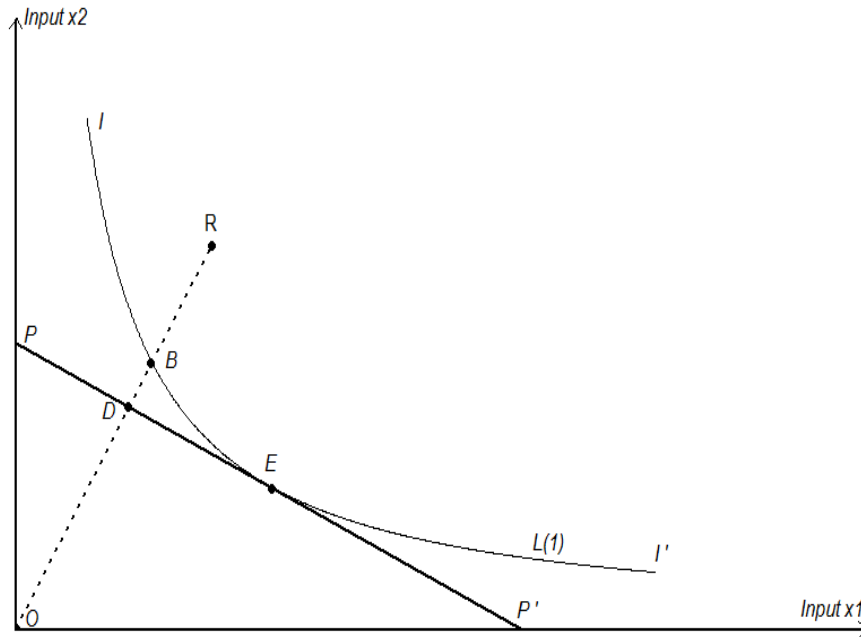
In Figure 1.3 consider the isoquant II' and the isocost line PP' that minimizes total cost of producing one unit of output. Let the point R be a vector of input quantities to produce a unit output belonging to the input correspondence image set $L(1)$.

Then, for a given input vector, Farrell defines the degree of technical efficiency (TE) as the ratio $\frac{OB}{OR}$, the allocative efficiency (AE) as $\frac{OD}{OB}$ and finally the overall productive efficiency (OPE) or the total economic efficiency as $\frac{OD}{OR}$. Note that the product of technical and allocative efficiencies provides the overall efficiency and all three measures of efficiency are between zero and one. The Farrell technical efficiency will be denoted θ in the nonparametric and TE in the parametric frontiers.

Furthermore, the distance DB represents the reduction in production cost that would occur if production were to occur at the allocatively (and technically) efficient point E . The distance BR can also be interpreted in terms of a cost

⁴Graph presented by Ian Crawford, Alexander Klemm and Helen Simpson in 2003 in their work entitled "Measuring public sector efficiency".

Figure 1.3: Technical and allocative efficiencies



reduction. The line PP' represents the input price ratio.

By definition, the technical efficiency reflects the ability of the firm to obtain maximal output from a given set of inputs; the allocative efficiency reflects the ability of the firm to use the inputs in optimal proportions, given their prices. These two efficiency components are combined to obtain the total economic efficiency.

A DMU is technically efficient when it is situated on the frontier, so when it realizes the maximal level of outputs using a given quantity of inputs; it is allocatively efficient when it minimizes the total costs of its production with choosing a social optimal level of this production. However, the technical efficiency is generally adopted in the frontier analysis because the DMU argues in term of the purpose to be reached. Once the frontier is determined, the distance between the observed point representing the behavior of the DMU and the frontier is measured. The smaller this distance, the smaller is also the inefficiency and hence bigger is the efficiency, which is the objective for any decision making unit.

In the following, a brief description of the parametric and the nonparametric approaches will be presented but let us start by pointing out that the parametric frontier suffers generally from the risk of misspecification of the model, the restrictive assumptions on the frontier function and from some problems in finite samples, as for the nonparametric frontier suffers from the complexity essentially when it is stochastic. For both of them a deterministic frontier and a stochastic frontier can be defined. The deterministic frontier considers that every deviation from the frontier is an inefficiency; in contrary, the stochastic frontier considers that a part of the deviation represents the statistical noise.

1.1.1 The deterministic and the stochastic nonparametric approaches

In nonparametric frontier analysis we distinguish between the input oriented and the output oriented framework. The aim of the input oriented one is to reduce proportionally the inputs to reach the frontier while keeping the output level unchanged. This framework is generally used when the DMU can control the inputs but it has not the capability to control the outputs as for example public firms. As for the output oriented framework and contrary to the input oriented one, the principal idea is to maximize the output until reaching the frontier while keeping the input level unchanged.

Besides, we distinguish also in the nonparametric frontier analysis between the deterministic and the stochastic approaches but the frequently used model is the deterministic one. It is mainly a linear problem and the linear programming approach is used to solve it. The stochastic one is more recent and may not always be easy to handle, thus there remain open research issues in this area. Their model representations are as follows:

a. The deterministic nonparametric frontier

As mentioned previously, the literature presents mainly two nonparametric methods to estimate the efficiency scores which are the Data Envelopment Analysis (DEA) and the Free Disposal Hull (FDH). The first one considers the convexity of the production set Ψ , but this is not the case for the second. Ψ is always estimated from a sample $\chi =$

$\{(x_i, y_i), i = 1, \dots, n\}$ drawn from the observed population. These methods have been widely used to estimate technical efficiency in a variety of domains such as industries and the public sector.

This kind of approach requires the definition of the production technology of n DMUs based on p inputs $x \in \mathbb{R}_+^p$ and q outputs $y \in \mathbb{R}_+^q$. The technology is represented by its production possibility set $\Psi = \{(x, y) \mid x \text{ can produce } y\}$ which means the set of all feasible input-output vectors, while the boundary of Ψ is defined by $\Psi^\partial = \{(x, y) \in \Psi \mid (\theta x, y) \notin \Psi, \forall 0 < \theta < 1, (x, \gamma y) \notin \Psi, \forall \gamma > 1\}$ as before and indicates the production or the technology frontier. From the boundary Ψ^∂ the Farrell input efficiency score and output efficiency score for a DMU operating at (x, y) are defined respectively as

$$\theta(x, y) = \inf \{\theta \mid (\theta x, y) \in \Psi\}, \quad (1.1.6)$$

$$\gamma(x, y) = \sup \{\gamma \mid (x, \gamma y) \in \Psi\}. \quad (1.1.7)$$

$\theta(x, y)$ is the radial contraction of inputs the DMU should achieve to be input-efficient in the sense that $(\theta(x, y)x, y)$ is a frontier point. Similarly, $\gamma(x, y)$ is the proportionate increase of output the DMU should achieve to be output-efficient in the sense that $(x, \gamma(x, y)y)$ is on the frontier. Both measures are positive such that $\theta(x, y) \leq 1$ and $\gamma(x, y) \geq 1$ and a value of 1 indicates an efficient DMU.

When it is needed, for example to change the limits of the efficiency in certain analyses, one can alternatively measure the input and output distance functions of Shephard developed in Shephard (1970). The Shephard input and output efficiency scores are simply defined as the inverse of their corresponding Farrell efficiency scores as $(\theta(x, y))^{-1}$ and $(\gamma(x, y))^{-1}$.

a1. The DEA formulation

The DEA formulation as a linear programming was introduced by Charnes et al. (1978). The input-oriented (IO) formulation of the Farrell measure, in the case where the boundary Ψ^∂ of the production set is a VRS, is expressed by the following linear programming

problem:

$$\begin{aligned} & \min \theta \\ & s.t. \begin{cases} Y^t \lambda \geq Y_0, \\ \theta X_0 - X^t \lambda \geq 0, \\ I_n^t \lambda = 1, \\ \theta \geq 0, \lambda \geq 0, \end{cases} \end{aligned} \quad (1.1.8)$$

where the first constraint forces the virtual DMU to produce at least the outputs $Y_0 \in \mathbb{R}_+^q$ of the studied DMU; the second constraint determines how much less input the DMU would need. Thus, it has to produce at least Y_0 using at most the efficient level of inputs θX_0 . The factor used to reduce the inputs is θ and this value is the efficiency score of the DMU. The weight coefficient $\lambda \in \mathbb{R}_+^n$ is a vector of percentages of all DMUs; $X_0 \in \mathbb{R}_+^p$ is the input vector of the studied DMU; $X^t \lambda$ and $Y^t \lambda$ are the input and output vectors for the analyzed producer given that X and Y are the (n, p) input matrix and (n, q) output matrix respectively; I_n is a vector of n elements all equal to one. The input-oriented efficiency score $\theta_{DEA-IO}(x, y)$ is hence estimated for each DMU.

In the same way, in the output-oriented (OO) formulation when the VRS case is also considered, to estimate $\gamma_{DEA-OO}(x, y)$ the expression of the following linear program is used:

$$\begin{aligned} & \max \gamma \\ & s.t. \begin{cases} \gamma Y_0 - Y^t \lambda \leq 0, \\ X^t \lambda \leq X_0, \\ I_n^t \lambda = 1, \\ \gamma \geq 0, \lambda \geq 0, \end{cases} \end{aligned} \quad (1.1.9)$$

a2. The FDH formulation

Since it is often difficult to find a good theoretical or empirical justification for postulating convex production sets in efficiency analysis, not restricting oneself to a convex technology seems an attractive

property. Estimation in this context was proposed in Afriat (1972) and for the same purpose, Deprins, Simar and Tulkens (1984) have proposed the FDH estimator which does not require the convexity of the technology. So, to have the FDH formulation, the DEA linear programming problems are modified to consider this property.

Indeed, both the input oriented and the output oriented formulation of the FDH method differ from their DEA counterparts just by the constraint on the vector λ . The constraint $\lambda \geq 0$ is replaced by $\lambda \in \{0, 1\}$. By construction, solving the two linear programs provides the estimates $\theta_{FDH-IO}(x, y)$ and $\gamma_{FDH-OO}(x, y)$ for each DMU. Of course, all the DMUs declared efficient by the DEA method are also declared efficient by the FDH one.

a3. The consistency of the DEA and the FDH estimates

Knowing that an estimator does not usually coincide with the true value, the quality of the estimators is often studied using a certain number of statistical tools such as the bias, variance, and the consistency. The purpose thus is to control the committed error when θ is estimated by $\hat{\theta}_n$. Given the assumption of the convexity of Ψ in the DEA method, the consistency and the rates of convergence of the efficiency estimates $\theta_{DEA}(x, y)$ and $\theta_{FDH}(x, y)$ can be compared in the input oriented case. Before making this, let us define the consistency of the estimator.

To study the consistency of the estimators the convergence in probability is adopted. Thus, an estimate $\hat{\theta}_n$ is weakly consistent if it converges in probability to the quantity of interest θ , noted $\hat{\theta}_n \xrightarrow{P} \theta$ as $n \rightarrow \infty$, i.e. if

$$\lim_{n \rightarrow \infty} P(|\hat{\theta}_n - \theta| < \epsilon) = 1, \quad \forall \epsilon > 0. \quad (1.1.10)$$

This means that when the sample size increases, the estimator $\hat{\theta}_n$ converges in probability to the true and unknown value θ . Generally, the convergence in probability at rate n^α of the estimator $\hat{\theta}_n$ of θ

is written $\hat{\theta}_n - \theta = o_P(n^{-\alpha})$ and its weaker form is $\hat{\theta}_n - \theta = O_P(n^{-\alpha})$. The parametric estimation problems yields estimators that converge in probability at the rate $n^{1/2}$ and are said to be root- n consistent. This provides a familiar benchmark to which the rates of convergence of other nonparametric estimators can be compared. Nevertheless, it is often difficult to prove the convergence of an estimate and to obtain its rate of convergence in a nonparametric approach. However, a proof of the consistency of the estimated efficiency scores in the multivariate case, where $p > 1$ and $q > 1$, with their rates of convergence in the input oriented case was developed in Kneip, Park and Simar (1998) for the DEA case and in Park, Simar and Weiner (2000) for the FDH case. They obtain:

$$\hat{\theta}_{DEA}(x, y) - \theta_{DEA}(x, y) = O_P(n^{-\frac{2}{p+q+1}}), \quad (1.1.11)$$

$$\hat{\theta}_{FDH}(x, y) - \theta_{FDH}(x, y) = O_P(n^{-\frac{1}{p+q}}), \quad (1.1.12)$$

hence, there is no $\alpha > 2/(p + q + 1)$ for the DEA case and no $\alpha > 1/(p+q)$ for the FDH case such that $\hat{\theta}(x, y) - \theta(x, y) = O_P(n^{-\alpha})$. The convergence rates $n^{-2/(p+q+1)}$ and $n^{-1/(p+q)}$ are both affected by p and q simultaneously rather than by p or q . This reflects the curse of dimensionality which is even worse for the multivariate case given that the enlargement of the dimension weakens the rate. Indeed, if the dimension $p + q$ increases, α decreases and then $n^{-\alpha}$ increases ($p + q \nearrow \Rightarrow \alpha \searrow \Rightarrow n^{-\alpha} \nearrow$) which is not wished given that the rate of convergence is strictly smaller than one and that it is low when it increases and gets closer to one. For more details about the consistency and about the approximation of the sampling distribution of the estimator when n is large, see the works mentioned previously Kneip, Park and Simar (1998) and Park, Simar and Weiner (2000) and also Daraio and Simar (2007).

Given that we have interest that the rate of convergence should be the fastest possible, the comparison of the two rates indicates that the convergence rate of the FDH estimate is lower than that of the

DEA and then the convergence rate of the DEA is better because it is the fastest. This advantage in favour of the DEA is due to the convexity of the production set Ψ . Furthermore, in comparison with the parametric estimation, when $p = q = 1$ the convergence rate of the DEA is also better given that $2/3$ is greater than $1/2$ and that of the FDH is the same.

a4. The nonparametric robust frontiers

As an alternative to the traditional FDH and DEA approaches some robust frontiers presenting a number of advantages as finding outliers in frontier settings have been developed. We refer to the works of Cazals et al. (2002), Daraio and Simar (2005), Aragon et al. (2005), Daouia and Simar (2005, 2007), Wheelock and Wilson (2008), Wilson (2011), Simar and Vanhems (2012), and Simar et al. (2012). We present here in the input orientation an overview of the robust order- m estimators and robust order- α estimators which have the properties that they do not suffer from the curse of dimensionality and the standard parametric root- n rate of convergence is achieved along with asymptotic normality. More details and examples are found in Simar and Wilson (2013). These estimators involve the concept of partial frontiers, as opposed to the traditional idea of a full frontier Ψ^∂ that envelops all the data.

a4.1. The robust order- m frontiers

The order- m frontiers have been proposed in Cazals et al. (2002). In the input oriented case, its basic idea is to define the expected minimum input achieved by any m DMUs chosen randomly from the population of DMUs producing at least output level y . If m goes to infinity, the estimation becomes identical to FDH estimation of the full frontier Ψ^∂ which envelopes all of the sample observations. Considering the formulation based on the probability of being dominated, the joint probability of $(X, Y) \in \mathbb{R}_+^{p+q}$ can be defined as $H_{XY}(x, y) = Pr(X \leq x \mid Y \geq y) \cdot Pr(Y \geq y) = F_{X|Y}(x \mid y) S_Y(y)$, where $F_{X|Y}(x \mid y)$ is the conditional distribution function of X and $S_Y(y)$ is the survivor function of

Y . Suppose a given output level y and consider m independent and identically distributed (i.i.d.) random variables X_i , $i = 1, \dots, m$ drawn from the conditional p -variate distribution function $F_{X|Y}(\cdot | y)$, and define the set

$$\Psi_m(y) = \bigcup_{j=1}^m \{(\tilde{x}, \tilde{y}) \in \mathbb{R}_+^{p+q} \mid \tilde{x} \geq X_j, \tilde{y} \geq y\} \quad (1.1.13)$$

This random set is the free disposal hull of m randomly-chosen DMUs that produce at least the level y of output. Then, for any x and the given y , the Farrell input efficiency score relative to the random set $\Psi_m(y)$ is $\tilde{\theta}_m(x, y) = \inf \{\theta \mid (\theta x, y) \in \Psi_m(y)\}$. Given that $\Psi_m(y)$ is random, this efficiency score $\tilde{\theta}_m(x, y)$ is also random. A realization of $\tilde{\theta}_m(x, y)$ is obtained by

$$\tilde{\theta}_m(x, y) = \min_{i=1, \dots, m} \left\{ \max_{j=1, \dots, p} \left(\frac{X_i^j}{x^j} \right) \right\}. \quad (1.1.14)$$

The order- m input efficiency score is defined as follows: For all y such that $S_Y(y) = Pr(Y \geq y) > 0$, the expected order- m input efficiency measure, referred to as the order- m input efficiency score, is given by

$$\theta_m(x, y) \equiv E\left(\tilde{\theta}_m(x, y) \mid Y \geq y\right). \quad (1.1.15)$$

Hence, in the input orientation the partial frontier can be defined as $\Psi_m^{\partial in} = \{(\tilde{x}, \tilde{y}) \mid \tilde{x} = x\theta_m(x, y), \tilde{y} = y, (x, y) \in \Psi\}$. The $\theta_m(x, y)$ estimator is

$$\hat{\theta}_{m,n}(x, y) \equiv \hat{E}\left(\tilde{\theta}_m(x, y) \mid Y \geq y\right). \quad (1.1.16)$$

The output order- m efficiency score can be defined along the same lines, see Cazals et al. (2002). Besides, a Monte-Carlo algorithm for approximating the expectation can be found in Daraio and Simar (2005).

a4.2. The robust order- α frontiers

The idea underlying order- α quantiles is similar to order- m partial frontiers. Instead of benchmarking units' performances against the full frontier, one can benchmark DMUs's performances against a quantile lying close to the full frontier. In the input-oriented framework, the order- α quantile frontiers benchmarks the unit operating at (x, y) against the input level not exceeded by $(1 - \alpha) \times 100$ percent of DMUs among the population of units producing output levels of at least y . The order- α input efficiency score is defined as follows: For all y such that $S_Y(y) > 0$ and for $\alpha \in (0, 1]$, the α -quantile input efficiency score for the unit operating at $(x, y) \in \Psi$, is given by

$$\theta_\alpha(x, y \mid \Psi) = \inf \{ \theta \mid F_{X|Y}(\theta x \mid y) > 1 - \alpha \}. \quad (1.1.17)$$

The only difference between the full frontier $\theta(x, y \mid \Psi)$ and $\theta_\alpha(x, y \mid \Psi)$ is that $F_{X|Y}(\theta x \mid y) > 0$ has been replaced with $F_{X|Y}(\theta x \mid y) > 1 - \alpha$. If $\theta_\alpha(x, y \mid \Psi) < 1$ ($>$), the unit at (x, y) can reduce (increase) its input usage to $\theta_\alpha(x, y \mid \Psi)x$ to become input-efficient at the level $\alpha \times 100\%$.

An estimator of $\theta_\alpha(x, y)$ is obtained by replacing the conditional distribution by its empirical counterpart, yielding

$$\hat{\theta}_{\alpha,n}(x, y) = \inf \left\{ \theta \mid \hat{F}_{X|Y,n}(\theta x \mid y) > 1 - \alpha \right\}. \quad (1.1.18)$$

Similarly, the estimator of the output α -quantile efficiency score for a unit operating at $(x, y) \in \Psi$ is given for all x such that $F_X(x) > 0$ and for $\alpha \in (0, 1]$ by

$$\hat{\gamma}_{\alpha,n}(x, y) = \sup \left\{ \gamma \mid \hat{S}_{Y|X,n}(\gamma y \mid x) > 1 - \alpha \right\}. \quad (1.1.19)$$

Daouia and Simar (2007) provides an algorithm for computation of these estimators.

b. The stochastic nonparametric frontier

The deterministic nonparametric approach with its DEA and FDH tools does not assume a particular functional form for the frontier and attributes all deviations from the frontier to inefficiency and hence completely ignores any stochastic noise in the data. In the literature, several studies tried to propose a stochastic semiparametric or nonparametric model but the problem is difficult because without some restrictions on the model a stochastic nonparametric model is not identified. The most part of these studies combine the DEA and the SFA or start from the SFA principal and replace the frontier function by a kernel regression or local maximum likelihood techniques.

For instance, recently, Kneip, Simar and Van Keilegom (2015) treat this nonparametric approach with unknown frontier and unknown variance of a normally distributed error. The authors propose a nonparametric method identifying and estimating both quantities simultaneously. They establish the consistency and rate of convergence of their estimators, verify and illustrate the performance of the estimators for small samples with simulations and apply their method to data of American electricity companies.

As another way to handle this approach, Kuosmanen and Kortelainen (2012) combines the DEA-type nonparametric frontier, which satisfies monotonicity and concavity, with the SFA-style stochastic homoskedastic composite error term in an encompassing semiparametric frontier model. Specifically, the authors assume that the observed data deviates from a nonparametric, DEA-style piecewise linear frontier production function due to a stochastic SFA-style composite error term, consisting of homoskedastic noise and inefficiency components. To estimate the model a new two-stage method called the Stochastic Non-smooth Envelopment of Data (StoNED) is proposed. In the first step, the shape of the frontier is estimated by applying the Convex Nonparametric Least Squares (CNLS) regression, which does not assume a priori any particular functional form for the regression function and in the second step, the method of moments and pseudolikelihood techniques are used to estimate the conditional ex-

pectations of inefficiency given the CNLS residuals.

In order to avoid the functional form error in modeling the mean of inefficiency function and the distributional assumptions on the components of the errors (the noise and the inefficiency), Tran and Tsionas (2009) proposes a simple semiparametric stochastic frontier model under the assumption that the mean of technical inefficiency depends on a set of covariates via an unknown but smooth function. The main objective of the authors is to estimate the frontier coefficients, measure technical inefficiency, and estimate the marginal effects of covariates on the mean of technical inefficiency. They derive hence a simple two-step semiparametric estimator that can be implemented using nonparametric kernel regression and least squares technique. Their method was illustrated using an empirical large panel data set of British manufacturing firms.

Furthermore, Kumbhakar et al. (2007) proposes a more general and flexible nonparametric approach for stochastic frontier models where the method is based on the conditional local maximum likelihood principle. The approach is nonparametric in the sense that the parameters of a given local polynomial model are localized with respect to the covariates of the model. They derive a practical computation of the local linear estimators of the individual efficiency scores for an observation (X_i, Y_i) as in Jondrow et al. (1982). Given that $Y = r(X) + v - u$, these estimates are obtained from $\hat{u}_i = \hat{E}[u(X_i) | \hat{\epsilon}(X_i), X = X_i]$, where v is the noise, u is the inefficiency, $\hat{\epsilon}(X_i) = Y_i - \hat{r}_0(X_i)$ and $\hat{r}_0(X_i)$ is the local intercept estimator of the frontier function. If the variables are measured in logs, which is usually the case, then the efficiency score for an individual i is measured by $\widehat{TE}_i = \exp\{-\hat{u}_i\}$ which is between 0 and 1. In their numerical illustrations, they chose a bandwidth which is adjusted for different scales of the variables and different sample sizes.

Within the framework of panel data, Henderson and Simar (2005) have exploited recent advances in kernel regression estimation of continuous and categorical data to estimate fully nonparametrically the frontier and time variant technical efficiency. In their approach, they allow firm and time ef-

fects to be smoothed as well as the continuous regressors which makes the estimation procedure fully nonparametric and thus only requires minimal restrictions on the technology.

Besides, the purpose of Simar (2003b) is to adapt some previous results, such as the Hall and Simar (2002) methodology, on detecting change points, to improve the performances of the classical nonparametric DEA/FDH estimators in the presence of noise. Hall and Simar (2002) have proposed a technique which allows to estimate a boundary point in the presence of noise in the univariate and the bivariate cases and Simar (2003b) has extended this to the multivariate frontier setup. The numerical illustrations proposed by the author were based on simulations and have showed that the procedure works well if the noise to signal ratio is not too large and that the procedure appears also to be robust to outliers.

Earlier even, Park and Simar (1994) and Park et al. (1998) have employed the kernel function and a variety of the maximum likelihood methods in some particular cases. They generally consider the semiparametric estimation of the stochastic frontier panel models under various assumptions on the joint distribution of the random firm effects and the regressors and on various dynamic specifications.

In the same context, nonparametric stochastic models for the frontier are analyzed in Kneip and Simar (1996). The authors use the panel data to avoid the distributional assumptions. So, almost no restriction is imposed on the structure of the model or of the inefficiencies but they pointed out that the estimation needs large values of time periods T and of firm N to obtain reliable estimates of the individual production functions and estimates of the frontier function.

Of course, several other estimators were proposed in the literature for nonparametric and semiparametric stochastic frontier models and most of them are established using kernel regression. We can also mention among them McAllister and McManus (1993), Fan et al. (1996) and Adams et al. (1999). However, in the nonparametric approach there is often an identification issue to be addressed before the estimation of the model.

1.1.2 The deterministic and the stochastic parametric approaches

In the parametric approach, the difference between the deterministic and the stochastic frontier consists in whether a random noise term is included or not. Of course, this difference influences the technical efficiency expressions.

a. The deterministic frontier

The early studies in the frontier analysis assume that the frontier is deterministic. In the input-output space, the frontier is deterministic if the model is expressed by $Y_i = F(X_i, \beta) \exp \{-u_i\}$ or simply by

$$y_i = f(x_i, \beta) - u_i, \quad i = 1, \dots, n, \quad (1.1.20)$$

where $y_i = \log(Y_i)$ is the observed output of the DMU_i , $x_i = \log(X_i)$ the inputs quantities used by the DMU_i , β is a vector of unknown parameters to be estimated, F and f are appropriate production functions and u_i a non-negative variable representing the degree of inefficiency. So, the model considers that all deviations from the frontier are inefficiencies. The parameters of the model were estimated using the linear programming approach as described in Aigner and Chu (1968):

$$\begin{cases} \min \sum_{i=1}^n u_i \\ u_i \geq 0. \end{cases} \quad (1.1.21)$$

An estimator of the Technical Efficiency (TE) of the DMU_i was proposed, it is defined by

$$\widehat{TE}_i = \exp \{-u_i\} = \frac{\exp \{y_i\}}{\exp \{f(x_i, \hat{\beta})\}} \quad (1.1.22)$$

b. The stochastic frontier

The classical model in the parametric stochastic frontier case is defined by $y_i = f(x_i, \beta) + v_i - u_i$, $i = 1, \dots, n$, where $\epsilon_i = v_i - u_i$, v_i is the stochastic term and the functional form of f is generally assumed to be a linear function. So, a part of the deviation from the frontier is considered

as a statistical noise and the other part is an inefficiency. The traditional estimation of the model is based on two additional hypotheses which are:

H1: x_i , v_i and u_i are independent;

H2: u_i and v_i are independent and identically distributed (i.i.d.), where $v_i \sim N(0, \sigma_V^2)$ and u_i follows a distribution defined on $\mathbb{R}^+$.

The technical efficiency in this case is estimated by $TE_i = E\left(\exp\{-u_i\} \mid \epsilon_i\right)$. This approach will be developed in detail later in Chapters 3 and 4 with empirical examples of the local government and drinking water in Morocco.

1.2 What is local government in Morocco?

To implement the decentralization process, the country has tried gradually, from the 1960s, to transfer certain responsibilities and certain public authorities of the central government to territorial entities called local authorities. This transfer of responsibilities was accompanied by a transfer of financial resources to face the expenses.

The various local entities are managed by an assembly or a council elected by the citizens. Every assembly or council elects a bureau composed of a president and a number of assistants which depends on the population size of the entity. At first level of the local government there are regions, each one is divided in provinces and the province in turn is composed of municipalities and rural districts.

1.3 The drinking water sector in Morocco

Being a sensitive sector, the production of drinking water is largely managed by the public institution named the National Office of the Drinking Water and called ONEP. The water supply is delegated to private operators in four agglomerations which are LYDEC in Casablanca, Redal in Rabat-Salé and Amendis in Tanger and Tetouan while it is insured by the self-governances in twelve others

and by the ONEP in the rest of the country. Hence, the ONEP sells drinking water either directly to subscribers or to self-governance and private operators to distribute it to the various consumers.

The main drinking water resources in Morocco are conventional resources such as surface water and groundwater but also the desalination of sea water has been developed. It is noted that recently the reuse of treated wastewater was adopted in the irrigation of the agricultural grounds and the municipal green spaces instead of the drinking water with the aim of saving it.

1.4 The data motivation

Empirical studies handled as illustrative examples in this research are about some development areas in Morocco. As mentioned before, several sectors are susceptible in a frontier analysis using its parametric or nonparametric approaches but we focus the interest on specially the finance of the local government and the drinking water sector in Morocco.

Measurement of local government efficiency is an important task in order to evaluate the management performance of the elected representatives chosen by the citizens and to compare the performance of the local councils. This area continues to gain in popularity recently in the efficiency analysis. In this framework, it seems important to analyze the financing of the Moroccan provinces using the available variables. Regrettably the panel data of this kind of local authorities finance is not easily available; it can be obtainable within the central administration in Rabat but requires some efforts, times and means to be collected. So, analysis will be done in the cross-section case for this kind of data.

Besides, it seems abnormal to satisfy itself with a simple descriptive analysis of certain data as I had the opportunity to collect during my professional experience with the local authorities. This experience had allowed me to understand the framework of the municipal councils.

As for the water sector, it was a good experimental example to treat the panel data regarding the dependence between the noise term and the inefficiency term and knowing that it is not a simple task to have panel data on the de-

velopment domains in Morocco. Furthermore, the limitation of water resources and the continuous increase of its needs placed the problem of the management of water resources and its cost among the most urgent priorities at the local and world level. Therefore, it is essential to plan and to organize the resources mobilization of available water and to preserve and use these resources in a rational way in the various domains.

In this research both nonparametric and parametric techniques will be used given that each one is useful to reach a determined purpose. Nevertheless, each of them has its advantages and its inconveniences.

1.5 Comparison between nonparametric and parametric approaches

Of course the nonparametric methods DEA and FDH and the parametric SFA are used by several authors in different areas to estimate technical efficiency in order to compare DMUs, however, each has its adepts. Hence, saying that either one is better than the other is not appropriate. Nevertheless, the comparison between the two approaches can be done according to the procedure adopted and the assumptions made by each of them. So, each one is more useful in particular situations and has advantages and limitations. The principal differences are about the way that the statistical noise and the number of the outputs are handled, but other possible differences exist and some of them will be summarized in this section. Besides, some resemblances will also be briefly addressed.

Next to DEA and FDH methods which ignore the noise and consider that any deviation from the frontier is inefficiency, statistical methods proposed in the literature of the frontier analysis such as the SFA take into account the statistical noise. The SFA constructs a parametric frontier which accounts for stochastic errors but requires specific assumptions about the technology and the inefficiency term which may be inappropriate or very restrictive such as the half-normal or the exponential inefficiency. Moreover, Ritter and Simar (1997) mention that even in a parametric framework, allowing for the noise in frontier

models, the identification problems can become serious if they imply ambiguity in the interpretation of the results.

Since DEA considers all random noise to represent inefficiency, it is likely that SFA may yield a higher efficiency than DEA and that DEA may estimate more ambitious targets than SFA because the efficiency rating simply measures the distance between the current and the target input-output levels.

On the other hand, DEA has the advantage that it is able to manage complex production environments with multiple input and output technologies. It assumes a linear function of p inputs producing q outputs and envelops p inputs and q outputs in a space of $p + q$ dimension. However, SFA handles models with just one output, or an a priori weighted average of multiple outputs.

As for the technology function, the DEA does not require a parametric specification of a functional form to define the frontier; it assumes a linear function of outputs and a linear function of inputs. Furthermore, being a nonparametric method, DEA does not produce the usual diagnostic tools with which to judge the goodness of fit of the model specifications. Being a parametric method, the SFA method assumes a particular functional form for the technology function which reflects the relationship between the output and the inputs used to produce it.

Besides, they differ also in the adopted estimation procedure. The DEA estimation procedure and the related efficiency estimates are based on a comparison of the input-output levels of a fixed individual. They can therefore be very sensitive to data swings at the fixed individual level. However, the SFA estimation procedure and its efficiency estimates are based on estimated parameter values in the regression model and are therefore not very sensitive to data swings at the individual level. Hence, the exactness and the stability of efficiency estimates are entirely in accordance with the nature of the two methods. Furthermore, according to their procedures, the DEA method focuses on individual DMUs, while the SFA method focuses on the estimation of the frontier.

Because the frontier is constructed using all data including extreme observations, Simar and Wilson (1998, 2000a) and Cazals et al. (2002) noted that the envelopment estimators are very sensitive to outliers and extreme values,

especially when data may be contaminated by measurement error. One way to resolve this problem is to detect them according to techniques proposed in this setting as for example in Wilson (1993), Simar (2003a) or in Tran et al. (2010) and then delete them from the database. One has proposed a way to define robust nonparametric estimators of the frontier which are particularly easy to compute and are useful for detecting outliers.

Another difference exists between both methods, it touches the returns to scale. Indeed, the DEA program, in contrary to SFA, includes the returns to scale equation in order to take into account the characteristic of the constant, increasing, decreasing or variable returns to scales of data.

According to the assumptions, it should be noted that generally the non-parametric approach DEA has the disadvantage to ignore the statistical noise, but being a nonparametric way it requires minimal assumptions on the production frontier. On the contrary, the parametric SFA models have the advantage of allowing for statistical noise, but they require strong assumptions about the inefficiency term. On the other hand, either DEA or SFA have some non-testable assumptions as the no measurement error in DEA and the particular error distribution in SFA. Furthermore, both methods may be vulnerable to measurement and misspecification error with the risk of omitting significant variables, the inclusion of irrelevant variables or the adoption in SFA of an inappropriate technology. Thus errors in specification and estimation can largely affect both techniques.

In the literature, several articles handled the comparison between both methods such as Bowlin et al. (1985), Thanassoulis (1993), Cubbin and Tzanidakis (1998), Linna and Häkkinen (1998), Dolton, Marcenaro and Navarro (2001), Sampaio, Barros and Ramajo (2005), Greene (2005) and others but the judgment remains still not definitive about when and why there is convergence between them.

Finally, we can notice that in comparison with DEA, the stochastic frontier approach is the most popular in the area of agriculture particularly. Moreover, DEA has a strong community in other areas such as management science, banking, health, and electricity and generally in industrial areas where there are multiple outputs. Hence, the approach based on the regression is the favorite

of applied econometricians and the DEA approach is that of the management scientists. We could also summarize the comparison by the idea that for the domains where the measurement errors are considered the SFA is more appropriate, while the DEA approach is more appropriate elsewhere, but both methods possess their respective strengths and weaknesses.

This research consists mainly of three works related to the estimation and the construction of the confidence intervals of the technical efficiencies in both nonparametric and parametric stochastic frontier analysis approaches. The first work, presented in chapter 2, is a nonparametric study which treats the Data Envelopment Analysis and its application to the financing of the Moroccan municipalities; the second and the third topics, provided respectively of Chapter 3 and Chapter 4, are parametric studies presenting the Maximum Likelihood Estimation (MLE), the technical efficiencies and two proposed procedures of inference in the cases of cross-sectional and panel data respectively when the two components of the error term are dependent. Chapters 2 and 3 were developed in El Mehdi and Hafner (2014a, 2014b).

So, our main contributions in the frontier analysis are to address a statistical inference which introduces the bivariate and the multivariate copula functions in the bootstrap procedure which serves this purpose. This copula density will model the association, if it exists, between the noise and the inefficiency terms of the frontier model. Indeed, some elliptical as the Gaussian copula, Archimedean as the Ali-Mikhail-Haq (AMH), Clayton and Frank copulas and other copula families as the Fairlie-Gumbel-Morgenstern (FGM) family are used. The principle is to model the inputs-output relationship under the dependence in the error term density to estimate efficiency and, in the bootstrap procedure of the confidence intervals for the chosen model and for each bootstrap replication, draw dependently the two components using a bivariate copula density for the cross-sectional case and draw the noise variables independently each other but dependently with the inefficiency variable for the panel data. Besides, the existence of the association is performed with the nonparametric Kendall's statistical test related to the copula case.

Another contribution of our research is to propose new models describing the behavior of the inefficiency function through the time in the panel data. These

sinusoidal models describe a possible periodic relationship of the inefficiency in time which tends generally to disappear at the end of period under study. In addition, in comparison with some models known in the frontier literature since the 90s, one of our proposed models will be selected as the best pattern describing the inefficiency of our data.

Finally, for all approaches, we will develop the methodology and the empirical illustration results of real data not handled before. We will conclude by prospects for future research, mainly introducing the dependence between the two error terms when panel data present a positive skewness. As computers have become faster, bootstrap and resampling methods became feasible for inference and made all these results available for our model framework.

Chapter 2

Local government efficiency: The case of Moroccan municipalities

2.1 Introduction

Analysis of local government efficiency has attracted considerable interest over recent years. In times of scarce public budgets and increasing public debt, the efficiency of local administration of decentralized budgets becomes important also for the central government. Examples of recent studies are Athanassopoulos and Triantis (1998) for Greece, Loikkanen and Susiluoto (2005) for Finland, Prieto and Zofio (2001) and Balaguer-Coll et al. (2007) for Spain, Worthington (2000) for Australia, and Afonso and Fernandes (2008) for Portugal. Afonso and Fernandes (2008) give an excellent account of the literature and almost all research studies invoke the concept of cost or production efficiency introduced by Farrell (1957). While there are many efficiency studies for African economies, see e.g. Ogundari et al. (2012) for the Nigerian agriculture and

Mlambo and Ncube (2011) for the South African banking sector, the literature on local government efficiency in African countries is very scarce. This work tries to fill this gap.

In most studies of local government efficiency, nonparametric approaches of frontier analysis are used, in particular data envelopment analysis (DEA) assuming convexity of the production set, and free disposal hull (FDH) without the convexity assumption. Parallel to applications, progress has been made in understanding the statistical properties of DEA and FDH efficiency estimates, which however has not yet been entirely taken into account in applied research. For example, the use of bias corrected estimators based e.g. on the bootstrap seems to be an important ingredient for reliable interpretation of estimated efficiency scores. In this work, we study the management of the financial resources of Moroccan municipalities. We use an aggregate output measure, the financial autonomy, which is economically meaningful and avoids the typical curse of dimensionality of nonparametric estimators in high dimensions. To the best of our knowledge, no results are available in the literature on local government efficiency in Morocco. We try to fill this gap, accounting for recent statistical research in DEA and FDH and comparing the results with those of other countries.

A general result of our study is that very few Moroccan municipalities are close to the frontier according to DEA. By construction, efficiency scores are higher for FDH, and more municipalities are close to the frontier. However, even for FDH, more than 90% of the municipalities are inefficient. Moreover, we find that there is a negative relation between population size and efficiency, both for DEA and FDH. This differs from the results of Balaguer-Coll et al. (2007) for Spain. The remainder of the chapter is organized as follows. Section 2 presents the DEA framework, Section 3 describes the statistical approach of using the bootstrap to do inference on the efficiency score estimates, Section 4 presents FDH as an alternative to DEA with less restrictive assumptions, and Section 5 develops the application to our Moroccan data set. Section 6 concludes.

2.2 The DEA program

Data envelopment analysis (DEA) is a nonparametric programming approach used to determine the production frontier and to estimate technical efficiencies of decision making units (DMU) such as firms or countries. The concept of technical efficiency as introduced by Farrell (1957) has been widely used in empirical research on production efficiency.

We focus on input-oriented DEA, i.e. the question by how much input quantities can be proportionally reduced without changing the output quantities. This is the typical problem for public decision makers which have to ensure public services while trying to minimize the inputs, see e.g. Daraio and Simar (2007, p.30).

To describe the analytical structure of DEA, the input matrix is denoted by $X_{(n,p)}$ and the output matrix as $Y_{(n,q)}$, where p and q are respectively the number of inputs and outputs, and n is the number of DMU's under study. We suppose that each DMU produces the same q outputs in possibly different amounts using the same p inputs also in possibly different amounts. As in Sueyoshi (1999), we assume that all DMUs have linearly independent input and output vectors in their data domain. The DEA matrix formulation for a given point (X_0, Y_0) in the case of variable returns to scale is given by the following linear programming primal problem, which needs to be solved n times, once for each DMU, see e.g. Banker et al. (2004),

$$\begin{aligned} & \min \theta \\ & s.t. \begin{cases} \sum_{i=1}^n \lambda_i y_{ki} \geq y_{k0}, & k = 1, \dots, q \\ \sum_{i=1}^n \lambda_i x_{ji} \leq \theta x_{j0}, & j = 1, \dots, p \\ \sum_{i=1}^n \lambda_i = 1, \\ \theta \geq 0, \lambda_i \geq 0 & \forall i = 1, \dots, n \end{cases} \end{aligned} \quad (2.2.1)$$

where θ corresponds to the level of technical efficiency, $Y_0 = (y_{10}, \dots, y_{q0}, \dots, y_{q0})$ and $X_0 = (x_{10}, \dots, x_{j0}, \dots, x_{p0})$ are levels of outputs and inputs of the considered DMU. The variables $\lambda_i, i = 1, \dots, n$ insure the convex hull of inputs and outputs in these data spaces. The restriction $\sum_{i=1}^n \lambda_i = 1$ corresponds to the assumption of variable returns to scale (VRS). It can be replaced by other

assumptions on the returns to scale (RTS), namely $\sum_{i=1}^n \lambda_i > 1$ (increasing returns to scale) and $\sum_{i=1}^n \lambda_i < 1$ (decreasing returns to scale). If the restriction on $\sum_{i=1}^n \lambda_i$ is dropped, one obtains constant returns to scale (CRS).

Let θ^* denote the optimal level of the efficiency score. In the case of one input and one output, θ_i is a radial measure of the distance between (x_i, y_i) and the corresponding frontier. The DMU is efficient when $\theta^* = 1$ and inefficient in case of $0 \leq \theta^* < 1$.

The slack variables S_k^- and S_j^+ associated with the dual variables u_k and v_j respectively lead us to the following program:

$$\begin{aligned} \min \quad & \left(\theta + \sum_{k=1}^q S_k^- + \sum_{j=1}^p S_j^+ \right) \\ \text{s.t.} \quad & \begin{cases} \sum_{i=1}^n \lambda_i y_{ki} - S_k^- = y_{k0}, & k = 1, \dots, q \\ \sum_{i=1}^n \lambda_i x_{ji} + S_j^+ = \theta x_{j0}, & j = 1, \dots, p \\ \sum_{i=1}^n \lambda_i = 1, \\ \theta \geq 0, \lambda_i, S_k^-, S_j^+ \geq 0 & \forall i = 1, \dots, n, \forall k, \forall j \end{cases} \end{aligned} \quad (2.2.2)$$

DMU_{i_0} is Pareto-efficient or fully efficient if and only if $\theta_{i_0}^* = 1$ and all slacks S_k^{-*} and S_j^{+*} are zero, see e.g. Thanassoulis (2001).

Instead of solving the primal program, it is often easier to use the dual program. In our case, the dual program is given by

$$\begin{aligned} \text{Max} \quad & \sum_{k=1}^q u_k y_{k0} + u^* \\ \text{s.t.} \quad & \begin{cases} \sum_{j=1}^p v_j x_{j0} \leq 1, \\ \sum_{k=1}^q u_k y_{ki} - \sum_{j=1}^p v_j x_{ji} + u^* \leq 0, & i = 1, \dots, n \\ u_k, v_j \geq 0 & \forall k, j \end{cases} \end{aligned} \quad (2.2.3)$$

where $U^t = (u_1, \dots, u_k, \dots, u_q) \in \mathbb{R}_+^q$ and $V^t = (v_1, \dots, v_j, \dots, v_p) \in \mathbb{R}_+^p$ are row vectors of the dual variables related to the constraints of the primal problem. Furthermore, a constraint with a strict equality in the primal is replaced by a free variable $u^* \in \mathbb{R}$ in the dual. Being free, this variable should be replaced by the difference between two positive variables t_1 and t_2 in a linear programming problem solved by the simplex method. See the Appendix A for more details

about how to get the dual program from the primal one.

At a point (x_0, y_0) we have $\hat{\theta}(x_0, y_0) = \sum_{k=1}^q \hat{u}_k y_{k0} + \hat{u}^*$.

The sign of the optimal value of u^* is used to identify the type of RTS at a point (x_0^*, y_0^*) on the efficiency frontier. Being negative, zero, positive or free, this variable indicates NIRS, CRS, IRS or VRS respectively.

It is useful to indicate that input and output oriented models may give different results with respect to their returns to scale. Thus, increasing returns to scale may result from an input oriented model, while an application of an output oriented model may produce a decreasing returns to scale for the same data. Also, it is worthwhile to note that working in smaller dimensions tends to provide better estimates of the frontier.

Many software packages include algorithms to solve linear programming problems of the type discussed previously. The problem can be cast in a form treatable by the simplex method used by R, see Appendix D for some commands.

2.3 Inference using the bootstrap

The bootstrap is a method which can be useful in many problems of statistical inference such as constructing confidence intervals and hypothesis tests. Its principle is to create a pseudo-replicate data set from the given data set, and then perform statistical inference using the replication set. The use of the bootstrap method in DEA goes back to Simar (1992).

The bootstrap method is based on the idea that the bootstrap distribution will mimic the original unknown sampling distribution of the estimators of interest (efficiencies) using a nonparametric estimate of their densities. Hence, a bootstrap procedure can simulate the data generating process (DGP) by using a Monte Carlo approximation and may provide a reasonable estimator of the true unknown DGP.

2.3.1 Data Generating Process

Consider a statistical model where a DGP P generates a random sample $\chi = \{(X_i, Y_i)_{i=1}^n\}$ of size n and suppose that we want to investigate the sampling distribution of the estimator $\hat{\theta}$ of an unknown parameter θ .

Using the nonparametric method described in (2.2.1) it is possible to estimate θ by $\hat{\theta}$ at a fixed point (x, y) for each DMU. As the DGP P is unknown, the bootstrap procedure is used to determine the DGP $\hat{P}$ as an estimator of the true unknown DGP. Thus, since $\hat{P}$ is known, we can generate a data set $\chi^* = \{(X_i^*, Y_i^*), i = 1, \dots, n\}$ from $\hat{P}$. This pseudo-sample defines the quantities $\hat{\theta}^*$ corresponding to the efficiencies $\hat{\theta}$ at the point (x, y) .

Analytically, it may be difficult to compute the true distribution of $\hat{\theta}^*(x, y)$ resulting from a sample χ^* drawn from $\hat{P}$. Therefore, the Monte Carlo approximation can be employed to obtain the sampling distribution of $\hat{\theta}^*(x, y)$. Using $\hat{P}$ to generate B pseudo-samples χ_b^* for $b = 1, \dots, B$ and applying the model (2.2.1), we obtain a set of pseudo estimates $\{\hat{\theta}_b^*(x, y)\}_{b=1}^B$. These pseudo estimates give an approximation of the unknown sampling distribution of the efficiency scores $\hat{\theta}_b^*(x, y)$ conditional on $\hat{P}$.

2.3.2 Bootstrap correcting bias for DEA efficiency scores

The bootstrap algorithm allows us to obtain bias corrected estimators and to make inference on the DEA efficiency scores. Correcting for the bias introduces additional noise and thus increases the variance of the estimator. However, Daraio and Simar (2007) suggests that a bias correction should be considered in almost all practical situations. However, before defining the bias corrected estimator of DEA efficiency scores, we define the bias and the standard deviation of this estimator at a point (x, y) .

First, denote the estimator at point (x, y) of DEA efficiency score $\theta(x, y)$ by $\hat{\theta}(x, y)$ and its bootstrap estimator by $\hat{\theta}^*(x, y)$.

These bias and standard deviation of $\hat{\theta}(x, y)$ cannot be computed because its sampling distribution is unavailable and its asymptotic approximation is too complicated to handle. Nevertheless, a bootstrap approximation is available

and given by

$$\widehat{bias}^*(\hat{\theta}(x, y)) \approx \frac{1}{B} \sum_{b=1}^B \hat{\theta}_b^*(x, y) - \hat{\theta}(x, y); \quad (2.3.1)$$

$$\widehat{std^2}^*(\hat{\theta}(x, y)) \approx \frac{1}{B} \sum_{b=1}^B \hat{\theta}_b^{*2}(x, y) - \left(\frac{1}{B} \sum_{b=1}^B \hat{\theta}_b^*(x, y) \right)^2; \quad (2.3.2)$$

Then, the bias corrected estimator is

$$\tilde{\theta}(x, y) = \hat{\theta}(x, y) - \widehat{bias}^*(\hat{\theta}(x, y)) = 2\hat{\theta}(x, y) - \frac{1}{B} \sum_{b=1}^B \hat{\theta}_b^*(x, y); \quad (2.3.3)$$

In (2.3.3), the correction is done by the mean. If the distribution of $\hat{\theta}^*(x, y)$ is asymmetric, the correction by the median can be used and may be more appropriate. In that case, we define the bias corrected estimator by

$$\tilde{\theta}(x, y) = 2\hat{\theta}(x, y) - \text{median}(\hat{\theta}_b^*(x, y), b = 1, \dots, B); \quad (2.3.4)$$

Generally the bias tends to disappear when the sample size is large. Gijbels et al. (1999) derive, when a single input and a single output are considered, the asymptotic distribution of the DEA estimator, based on which they propose an improved bias-corrected estimator.

2.3.3 Confidence intervals for DEA efficiency scores

To do statistical inference and, in particular, to construct confidence intervals we need the distribution function of the variable of interest for computing or estimating the quantiles. Since in DEA the sampling distribution of $W = \hat{\theta}(x, y) - \theta(x, y)$ is unknown, the bootstrap method will provide an appropriate approximation, see e.g. Daraio and Simar (2007).

By definition, the efficiency's confidence interval at level $1 - \alpha$, for all $\alpha \in$

$[0, 1]$, is

$$P\left(\hat{\theta}(x, y) - a_{1-\frac{\alpha}{2}} \leq \theta(x, y) \leq \hat{\theta}(x, y) - a_{\frac{\alpha}{2}}\right) = 1 - \alpha; \quad (2.3.5)$$

The method adopted to build the bootstrap confidence interval for efficiency is the basic bootstrap method that adjusts automatically for the bias of the DEA estimator. The bootstrap approximation of the confidence interval for $\theta(x, y)$ is given by

$$P\left(\hat{\theta}(x, y) - \hat{a}_{1-\frac{\alpha}{2}} \leq \theta(x, y) \leq \hat{\theta}(x, y) - \hat{a}_{\frac{\alpha}{2}}\right) \approx 1 - \alpha; \quad (2.3.6)$$

where $\hat{a}_{\alpha'} = \hat{c}_{\alpha'} - \hat{\theta}(x, y)$ and $\hat{c}_{\alpha'}$ is the α' -quantile of the empirical distribution of the estimators $\left\{\hat{\theta}_b^*(x, y)\right\}_{b=1}^B$.

As usual, the precision will be higher when the DEA frontier above (x, y) is determined by many sample points (X_i, Y_i) , as the length of the confidence interval will be smaller, and vice versa. Simar and Wilson (2000a, 2000b) have shown that the naive bootstrap described above is inconsistent, but a smoothed version of it can be shown to be consistent. The smoothed bootstrap of **FEAR** (the Frontier Efficiency Analysis with R) developed by Wilson (2008) can be used to generate bootstrap replications, and this method is statistically consistent. The problem of the inconsistency of the naive bootstrap comes from the fact that the efficient facet that determines the value of $\hat{\theta}$ in the original sample χ appears too often, and with a fixed probability, in the pseudo-samples χ_b^* and this fixed probability does not vanish even when $n \rightarrow \infty$.

2.3.4 Testing returns to scale

The bootstrap can also be used for hypothesis tests, e.g. testing the returns to scale. The least restrictive model for returns to scale is the varying returns to scale (VRS) situation where the returns to scale are allowed to be locally increasing, then constant and finally non-increasing. Testing returns to scale (RTS) is carried out according to the following procedure where we test CRS against VRS, see e.g. Simar and Wilson (2002) and Daraio and Simar (2007).

Let Ψ be the production set, defined by

$$\Psi = \{(x, y) \in \mathbb{R}_+^{p+q} / x \in \mathbb{R}_+^p, y \in \mathbb{R}_+^q, (x, y) \text{ feasible}\}$$

To test the null hypothesis $H_0 : \Psi^\partial \text{ is CRS}$ against the alternative $H_1 : \Psi^\partial \text{ is VRS}$, we first estimate the efficiency scores at all points (X_i, Y_i) for the two cases CRS and VRS, denoted respectively $\hat{\theta}_{CRS}(X_i, Y_i)$ and $\hat{\theta}_{VRS}(X_i, Y_i)$. Then we define the test statistic

$$T(\chi_n) = \frac{1}{n} \sum_{i=1}^n \frac{\hat{\theta}_{CRS,n}(X_i, Y_i)}{\hat{\theta}_{VRS,n}(X_i, Y_i)}. \quad (2.3.7)$$

Under H_0 , $T(\chi_n)$ will be close to one since $\hat{\theta}_{CRS,n}(X_i, Y_i)$ and $\hat{\theta}_{VRS,n}(X_i, Y_i)$ are close to each other. By construction, $\hat{\theta}_{CRS,n}(X_i, Y_i) \leq \hat{\theta}_{VRS,n}(X_i, Y_i)$, and hence, under the alternative, $T(\chi_n)$ will be close to zero. Therefore, we reject H_0 for small values of $T(\chi_n)$, or formally, at level $\alpha \in (0, 1)$ if $p\text{-value} < \alpha$, where

$$p\text{-value} = P(T(\chi_n) < T_{obs} \mid H_0 \text{ is True}) \quad (2.3.8)$$

and T_{obs} is the value of $T(\chi_n)$ computed with the original observed sample χ_n . This probability cannot be computed analytically but we can approximate it by using the bootstrap by

$$p\text{-value} = \frac{1}{B} \sum_{i=1}^B I(T^{*b} \leq T_{obs}), \quad (2.3.9)$$

where $T^{*b} = T(\chi_n^{*b})$ is the value of T computed for each bootstrap sample, B is the number of pseudo samples χ_n^{*b} , and $I(\cdot)$ is the indicator function.

2.4 The free disposal hull approach

The DEA approach is based on a restrictive convexity assumption on the structure of the production set Ψ . Deprins, Simar and Tulkens (1984) have proposed an estimator supposing that the frontier of the production set is simply the

boundary of the free disposal hull (FDH) of the data set. The FDH approach to estimate the frontier only requires strong disposability of inputs and outputs and variable returns to scale, not convexity. A DMU is inefficient in the FDH sense if it is dominated by at least another DMU.

In the input oriented case, the FDH efficiencies at a fixed point (x_0, y_0) , denoted $\hat{\theta}_{FDH}(x_0, y_0)$, can be estimated by solving the following linear program that has $\lambda_i \in \{0, 1\}$ instead of $\lambda_i \geq 0$ in comparison with the DEA linear program:

$$\begin{aligned} \min \quad & \theta \\ \text{s.t.} \quad & \begin{cases} \sum_{i=1}^n \lambda_i y_{ki} \geq y_0, & k = 1, \dots, q \\ \sum_{i=1}^n \lambda_i x_{ji} \leq \theta x_0, & j = 1, \dots, p \\ \sum_{i=1}^n \lambda_i = 1, \\ \lambda_i \in \{0, 1\} & \forall i = 1, \dots, n \end{cases} \end{aligned} \quad (2.4.1)$$

In applications, these scores can be calculated as follows. For the sample $\chi = \{(X_i, Y_i), i = 1, \dots, n\}$ where $X_i \in \mathbb{R}_+^p$ and $Y_i \in \mathbb{R}_+^q$, let D_0 be the set of observations which dominate (x_0, y_0) ,

$$D_0 = \{i / (X_i, Y_i) \in \chi, X_i \leq x_0, Y_i \geq y_0\}.$$

Then,

$$\hat{\theta}_{FDH}(x_0, y_0) = \min_{i \in D_0} \left\{ \max_{j=1, \dots, p} \left(\frac{X_i^j}{x_0^j} \right) \right\}, \quad (2.4.2)$$

where X_i^j is the j^{th} component of $X_i \in \mathbb{R}_+^p$ and x_0^j is the j^{th} component of $x_0 \in \mathbb{R}_+^p$.

First, the maximum part of the algorithm identifies the dominant DMU's relative to which a given DMU is evaluated. Then, the estimators of the FDH efficiency scores are calculated from the minimum part of the algorithm. For each DMU declared inefficient by the FDH approach, it is possible to find at least one DMU in the set D_0 that presents a superior performance relative to the first DMU.

Simar and Wilson (2000a) have established the statistical properties of the

FDH estimator in a multivariate context, in order to do inference either by using asymptotic distributions or by means of bootstrap. The FDH estimator, like other nonparametric estimators such as DEA, suffers from the curse of dimensionality due its slow convergence rate in high dimensions. In our application we will reduce the dimension of input and output space to avoid this problem.

2.5 Illustrative application to local government efficiency in Morocco

To reduce the monopoly of the central administration in decision making, the kingdom of Morocco opted, since the first years of independence, for a system of decentralization. This system allows to involve the citizens with the management of local business and to give a sense of responsibility to the local leaders. Since the 1960's, the country tried progressively to transfer certain responsibilities and certain authorities of the central government towards well-defined local authorities. This transfer of responsibilities was accompanied by a transfer of financial resources to confront expenses.

Since 1997, the Moroccan local authorities include 16 regions, 68 prefectures and provinces and 1546 districts, of which 248 are urban and 1298 rural. The different local entities are managed by a council elected for a period of six years. Their financial transactions are established according to rules defined by the legislator and put back in a document called the budgetary document which describes the budget of the entity.

The budget plan of the local authority is prepared, approved and executed according to the current laws, regulations and instructions. Nevertheless, local authorities have the possibility of establishing secondary budgets for specific operations.

The main budget contains two parts. The first one describes the operating budget, and the second is the budget of equipment or investment. Each of them contains two parts, one describes receipts and the other one the expenses. In this framework, to facilitate the statistical analysis of the budgets of local authorities, the various budgetary columns were numbered according to a well

defined nomenclature. The budget is then divided in its two parts of recipes and expenses in sections, chapters, articles and paragraphs.

2.5.1 The data

We estimate the efficiency of the Moroccan rural districts by giving a particular attention to those of the oriental region. This region contains 91 rural districts. The inputs are constituted by ten variables which represent the categories of the financial resources of the local authority during the budgetary year 1998/1999. After the decentralization of the Moroccan administration, this kind of data is not available after 1999. Moreover, we were not able to obtain data before 1998, and therefore our analysis is restricted to only one budgetary year. We are aware that this may not be representative for an efficiency analysis of Moroccan municipalities, but due to the unavailability of the same variables in other periods, the analysis has to be restricted to this data set. All conclusions drawn from this analysis have therefore to be taken with care.

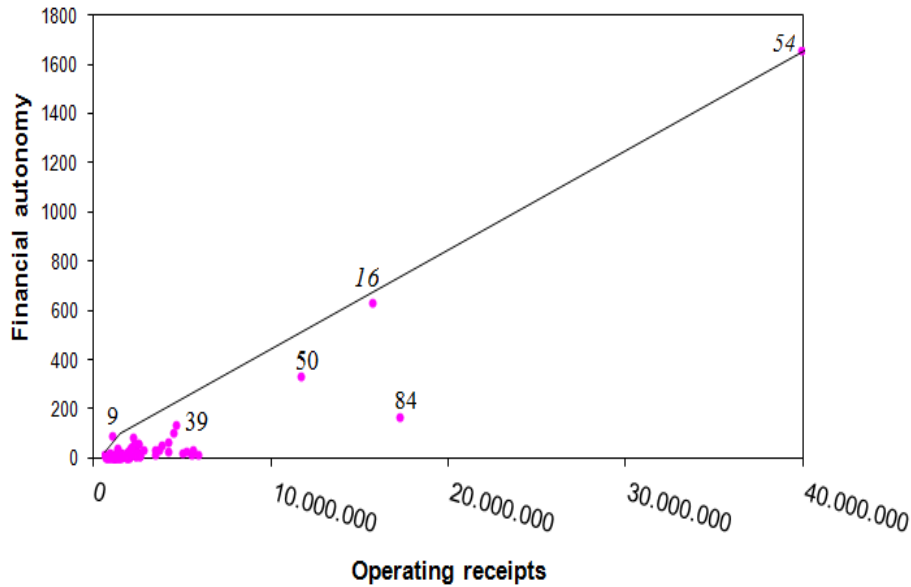
The ten variables describing the inputs are: The urban tax, the tax on the collection of the waste, the tax of the licence, the product of the forest domain, the taxes and assimilated taxes, the product of services, the product and the income of the goods, the concessions, the subsidies and competition and finally the order receipts. In order to reduce the dimension of the input space and thus to avoid the curse of dimensionality of nonparametric estimation, we decide to aggregate these ten input variables into a single one. See e.g. Daraio and Simar (2007, p.148) for a general justification of aggregate input and/or output measures. In our case, all input variables have the same scale and their unweighted sum has an economic meaning, which we call operating receipts. In addition, our aggregated variable is highly correlated with five of the ten inputs which are the assimilated taxes, the tax on the collection of the waste, the urban tax, the tax of the licence and the subsidies with a correlation rate of 0.942, 0.935, 0.931, 0.924 and 0.885 respectively. It is noted that the operating receipts less the subsidies are called the own receipts of the municipality.

With respect to output, we define a variable which measures the financial autonomy of the municipalities, defined as the ratio of the own receipts of the municipality and its operating expenses. If this ratio is one or larger, then the

municipality is financially autonomous, but not necessarily efficient. If the ratio is smaller than one, than it is not financially autonomous. Thus, we consider DEA and FDH efficiency estimates with a single input variable and a single output variable.

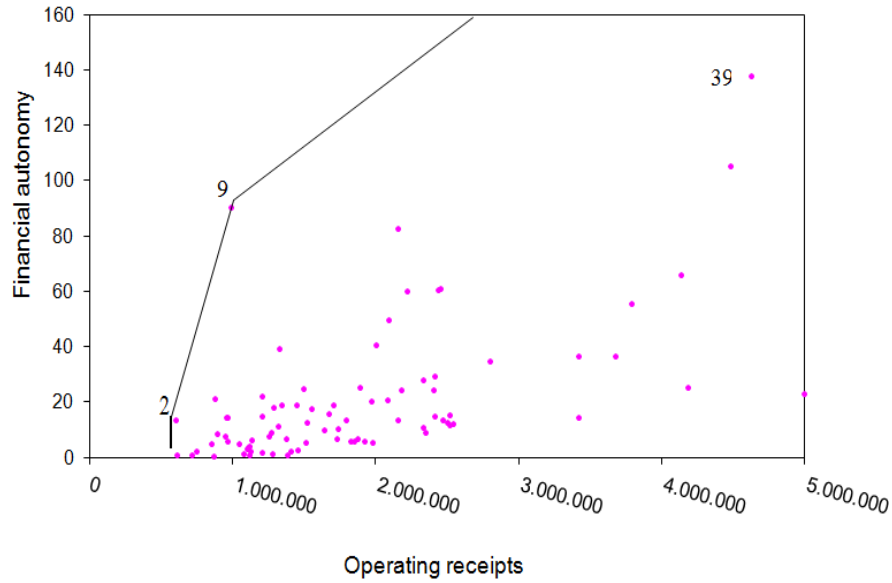
The data used to illustrate the approach consist of pairs (X_i, Y_i) where X_i represents the input expressed by the operating receipts of the DMU_i used to produce the output Y_i expressed by the financial autonomy for the same DMU_i . The relationship between the output and the input reveals a possible positive trend, as with higher operating receipts the financial autonomy increases, as shown in Figures 2.1 and 2.2. These figures also suggest a possible existence of outliers in our data set.

Figure 2.1: Financial autonomy versus operating receipts and DEA frontier



We see from Figures 2.1 and 2.2 that the districts Ras Asfour (2), Labkhata (9) and Ain Lehjer (54) are estimated as efficient (refer to Table 2.1 for the numbers of the districts). However, Laaouinate (16) is almost efficient because it is very close to the frontier. The Ain Lehjer (54) district is isolated from the others, so it may be possible that it is an outlier. If it dominates several districts, removing it with other outliers may declare some districts efficient

Figure 2.2: Financial autonomy versus operating receipts and DEA frontier (Zoom)



which were previously inefficient such as Tiouli (39). This finding confirms that frontier analysis is sensitive to outliers.

Since the presence of outliers may influence the efficiency scores, we used the procedure of Wilson (1993) of detecting outliers in deterministic nonparametric frontier models, one can use the robust order- m frontiers of Cazals et al. (2002). If outliers are identified, they will be deleted from the data set and efficiency scores will be re-estimated. The number of outliers being arbitrary in Wilson (1993), we first set it equal to ten and the procedure indicates that there are four possible outliers which are the districts of Laaouinate (16), Iksane (50), Ain Lehjer (54) and Selouane (84). Results show that deleting these observations from the data set does not influence substantially the efficiency scores of the other DMU's. Thus, we keep these districts in the data set.

2.5.2 Interpretation of the results

In the following, we present the estimation results and their interpretation, first for DEA and then for FDH.

DEA results

Before applying DEA to our data, we performed a hypothesis test about the returns to scale (RTS) in order to decide which DEA linear program to adopt. The statistical test described in Section 2.3 reveals the existence of variable returns to scale (VRS). Therefore, the linear program used to estimate the scores of efficiencies of DEA is that represented in the envelopment model expressed in (2.2.1) or the Multiplier model expressed in (2.2.3). To avoid the inconsistency of the naive bootstrap, we used the smoothed bootstrap of **FEAR** (the Frontier Efficiency Analysis with R).

Table 2.1: Efficiency scores of FDH, DEA with VRS and their confidence intervals

Pop Ord	District	Effc FDH	Domin FDH	Effc DEA	Mean Boot Eff DEA	Bias corr. Eff DEA	Bias DEA	Std DEA	Lower Bound CI	Upper Bound CI
1	2	3	4	5	6	7	8	9	10	11
1	AIN-CHOUATER	0.4798	10	0.4798	0.5132	0.4511	0.0287	0.026	0.4085	0.4777
2	RAS ASFOUR	1	1	1	1.0934	0.9217	0.0783	0.0583	0.8395	0.9878
3	ABBOU LAKHAL	0.8397	2	0.8397	0.8907	0.7955	0.0442	0.0458	0.7156	0.8375
4	AL BARKANYENE	0.9899	2	0.9899	1.0497	0.9381	0.0518	0.0541	0.8436	0.9877
5	EL ATEF	0.6365	4	0.6365	0.6808	0.5983	0.0382	0.0345	0.5419	0.6337
6	MRIJA	0.9079	3	0.63	0.6874	0.5818	0.0482	0.0367	0.5285	0.6234
7	LABSARA	0.6913	6	0.6913	0.7331	0.6551	0.0362	0.0378	0.5891	0.6897
8	GAFAIT	0.5242	3	0.3501	0.386	0.3209	0.0292	0.0247	0.2875	0.3459
9	LEBKHATA	1	1	1	1.7615	0.7117	0.2883	0.4077	0.6168	0.936
10	OULAD SIDI ABDELHAKEM	0.4981	21	0.4981	0.5284	0.4719	0.0262	0.0271	0.4245	0.4968
11	OULAD M'HAMMED	0.8053	2	0.8053	0.8547	0.7625	0.0428	0.0438	0.6863	0.8031
12	SIDI BOUBKER	1	1	0.7345	0.8026	0.6779	0.0566	0.0457	0.613	0.7265
13	SIDI MOUSSA LEMHAYA	0.7035	2	0.7035	0.7482	0.6649	0.0386	0.0381	0.5995	0.7014
14	TAFUUGHALT	0.3476	26	0.3476	0.3709	0.3275	0.0201	0.0188	0.2962	0.3463
15	SIDI BOULENOUAR	0.5591	13	0.5591	0.5931	0.5297	0.0294	0.0305	0.4765	0.5577
16	LAAOUNATE	1	1	0.9336	1.895	0.6465	0.2871	0.7433	0.5336	0.8842
17	OULAD DAOUD ZKHANINE	0.9161	2	0.6358	0.6936	0.5871	0.0487	0.037	0.5333	0.6291
18	AFSOU	0.6772	3	0.6772	0.7263	0.6349	0.0423	0.0368	0.5757	0.6738
19	M'HAJER	0.4378	18	0.4378	0.4671	0.4125	0.0253	0.0237	0.373	0.4361
20	BNI MATHAR	0.8243	2	0.5379	0.5888	0.4958	0.0421	0.0344	0.4476	0.5313
21	MESTFERKI	0.5312	9	0.5312	0.5668	0.5006	0.0306	0.0287	0.4527	0.5292
22	GUENFOUDA	0.443	7	0.3228	0.3522	0.2983	0.0245	0.0197	0.2697	0.3195
23	OULAD GHIZIYEL	0.6259	6	0.6259	0.6678	0.5897	0.0362	0.0339	0.5333	0.6235
24	AIN SFA	0.3664	17	0.3664	0.3942	0.3426	0.0238	0.0201	0.3107	0.3645
25	GTETER	0.6044	6	0.4335	0.4721	0.401	0.0325	0.0256	0.3635	0.4287

Table 2.1: Efficiency scores of FDH, DEA with VRS and their confidence intervals

Pop Ord	District	3	4	5	6	7	8	9	10	11
26	SIDI BOUHRIA	0.5741	10	0.5741	0.6106	0.5426	0.0315	0.0311	0.4892	0.5724
27	RISLANE	1	1	0.3039	0.4954	0.2234	0.0805	0.1096	0.1904	0.2865
28	TAZAGHINE	0.4745	8	0.4745	0.5104	0.4436	0.0309	0.026	0.4023	0.4721
29	MESTEGMER	0.428	26	0.428	0.4543	0.4053	0.0227	0.0233	0.3648	0.4268
30	AZLAF	0.313	34	0.313	0.3333	0.2954	0.0176	0.0169	0.2667	0.3119
31	BOUMERIEME	0.7459	2	0.5515	0.6441	0.4856	0.0659	0.0663	0.4183	0.5414
32	OUARDANA	0.5454	11	0.5454	0.5794	0.5160	0.0294	0.0295	0.4647	0.5438
33	MELG EL OUIDANE	0.5434	12	0.5434	0.5772	0.5141	0.0293	0.0294	0.4631	0.5419
34	TALILIT	0.5354	15	0.5354	0.5683	0.5070	0.0284	0.0291	0.4563	0.534
35	OULAD AMGHAR	0.5375	18	0.5375	0.5700	0.5094	0.0281	0.0294	0.4581	0.5363
36	FEZOUANE	0.4591	2	0.4403	0.7151	0.3244	0.1159	0.1598	0.2774	0.4241
37	BNI KHALED	0.4356	33	0.4356	0.4620	0.4128	0.0228	0.0238	0.3712	0.4346
38	BNI MARGHINE	0.5366	11	0.5366	0.5706	0.5071	0.0295	0.029	0.4573	0.535
39	TIOULI	1	1	0.4661	0.7370	0.3486	0.1175	0.1674	0.2928	0.4396
40	TANCHERFI	0.6635	3	0.4397	0.4834	0.4040	0.0357	0.03	0.3627	0.4355
41	OULAD BOUBKER	0.7259	4	0.508	0.5534	0.4697	0.0383	0.0296	0.426	0.5025
42	HASSI-BERKANE	0.5619	9	0.3997	0.4351	0.3699	0.0298	0.0234	0.3353	0.3953
43	MAATARKA	0.6784	4	0.4826	0.5253	0.4465	0.0361	0.0282	0.4048	0.4772
44	AIT-MAIT	0.456	9	0.456	0.4941	0.4236	0.0324	0.0257	0.3849	0.4529
45	BOUCHAOUENE	0.6505	5	0.4703	0.5125	0.4348	0.0355	0.0281	0.3934	0.4651
46	BNI-SIDEL-LOUTA	0.3614	24	0.2529	0.2755	0.2338	0.0191	0.0147	0.2121	0.2501
47	BNI GUIL	0.3037	36	0.3037	0.3235	0.2867	0.017	0.0164	0.2588	0.3027
48	BNI OUKIL OULAD MHAND	0.4715	24	0.4715	0.5001	0.4466	0.0249	0.0257	0.4018	0.4703
49	MECHRAA HAMMADI	0.3351	16	0.3351	0.3664	0.3088	0.0263	0.0195	0.2813	0.331
50	IKSANE	1	1	0.6141	1.0878	0.4437	0.1704	0.3533	0.3639	0.5849

Table 2.1: Efficiency scores of FDH, DEA with VRS and their confidence intervals

Pop Ord	District	3	4	5	6	7	8	9	10	11
51	AMEJJAOU	0.4132	28	0.4132	0.4385	0.3912	0.022	0.0224	0.3521	0.412
52	ISLY	0.4944	2	0.3681	0.4323	0.3228	0.0453	0.0458	0.2772	0.361
53	SIDI LAHSEN	0.4731	2	0.374	0.4622	0.3177	0.0563	0.0636	0.2679	0.3665
54	AIN LEHJER	1	1	1	2.9890	0.6336	0.3664	1.7052	0.5321	0.9258
55	TIZTOUTINE	0.5122	8	0.3673	0.4000	0.3398	0.0275	0.0217	0.308	0.3633
56	TSAFT	0.4185	9	0.305	0.3327	0.2817	0.0233	0.0186	0.2547	0.3018
57	BNI-SIDEL-JBEL	0.2573	29	0.2573	0.2778	0.2398	0.0175	0.0143	0.2177	0.2557
58	BOUANANE	0.2398	34	0.2398	0.2609	0.2219	0.0179	0.0138	0.2018	0.2378
59	IFERNI	0.3972	21	0.3972	0.4230	0.3749	0.0223	0.0215	0.3385	0.3959
60	RAS-EL-MA	0.4241	7	0.2876	0.3193	0.2623	0.0253	0.0217	0.2339	0.2836
61	BOUDINAR	0.3212	30	0.3213	0.3428	0.3027	0.0186	0.0174	0.2737	0.32
62	TENDRARA	0.4065	4	0.3442	0.4539	0.2813	0.0629	0.0793	0.2351	0.338
63	DAR-EL-KEBDANI	0.4107	7	0.2827	0.3161	0.2565	0.0262	0.0231	0.2274	0.2788
64	TAFERSIT	0.2362	36	0.2362	0.2570	0.2186	0.0176	0.0136	0.1988	0.2342
65	TROUGOUT	0.3457	18	0.3457	0.3732	0.3222	0.0235	0.0192	0.2925	0.3436
66	SIDI ALI BELQUASSEM	0.3957	13	0.3957	0.4305	0.3662	0.0295	0.0227	0.333	0.3924
67	AIN ZOIRA	0.2785	22	0.2785	0.3045	0.2567	0.0218	0.0162	0.2338	0.2751
68	IAAZZANENE	0.4115	11	0.2727	0.2998	0.2505	0.0222	0.0186	0.2249	0.2701
69	BNI-TADJITE	0.3466	27	0.2425	0.2642	0.2243	0.0182	0.0141	0.2034	0.2399
70	AGHBAL	0.2432	32	0.2432	0.2659	0.2242	0.019	0.0142	0.2042	0.2403
71	TALSINT	0.455	8	0.3016	0.3316	0.277	0.0246	0.0206	0.2487	0.2987
72	AHL OUAD ZA	0.2614	6	0.2146	0.2744	0.1787	0.0359	0.0436	0.1496	0.2109
73	CHOUHIYA	0.5224	11	0.3655	0.3982	0.338	0.0275	0.0213	0.3065	0.3616
74	MADAGH	0.2556	33	0.2556	0.275	0.239	0.0166	0.014	0.2167	0.2543
75	IJERMAOUAS	0.3292	30	0.3292	0.3506	0.3107	0.0185	0.0178	0.2805	0.3281

Table 2.1: Efficiency scores of FDH, DEA with VRS and their confidence intervals

Pop Ord	District	3	4	5	6	7	8	9	10	11
		Effc FDH	Domin FDH	Effc DEA	Mean Boot Eff DEA	Bias corr. Eff DEA	Bias DEA	Std DEA	Lower Bound CI	Upper Bound CI
76	BNI-BOUFPLOUR	0.1581	34	0.1116	0.1215	0.1032	0.0084	0.0065	0.0936	0.1103
77	AHL ANGAD	0.403	3	0.3433	0.4568	0.2791	0.0642	0.0806	0.2336	0.3355
78	LAATAMNA	0.3249	29	0.3249	0.3467	0.3062	0.0187	0.0176	0.2769	0.3237
79	TEMSAMANE	0.3534	9	0.2523	0.2876	0.2258	0.0265	0.0252	0.1973	0.2486
80	MIDAR	0.2365	17	0.1579	0.1741	0.1448	0.0131	0.0111	0.1297	0.156
81	IHADDADENE	0.2695	10	0.1951	0.2243	0.1737	0.0214	0.0211	0.1508	0.1926
82	BEN TAIEB	0.2895	9	0.2096	0.241	0.1866	0.023	0.0227	0.162	0.2069
83	FARKHANA	0.1779	15	0.1288	0.148	0.1146	0.0142	0.0139	0.0995	0.1271
84	SELOUANE	0.6751	3	0.1674	0.2688	0.1246	0.0428	0.0653	0.1044	0.158
85	AREKMANE	0.1024	47	0.1024	0.1119	0.0944	0.008	0.006	0.086	0.1011
86	OULAD SETTOUT	0.2394	3	0.21	0.2935	0.1661	0.0439	0.0543	0.1402	0.2036
87	BOUARG	0.2553	31	0.1772	0.1933	0.1636	0.0136	0.0103	0.1486	0.1753
88	BOUGHRIBA	0.2384	37	0.2384	0.2583	0.2214	0.017	0.0134	0.2012	0.2368
89	BNI-CHIKER	0.1981	24	0.1293	0.1415	0.1192	0.0101	0.0083	0.1076	0.1277
90	ZEGZEL	0.4453	3	0.377	0.4972	0.3081	0.0689	0.0869	0.2575	0.3703
91	DRIOUCH	0.1916	16	0.1329	0.1491	0.1202	0.0127	0.0113	0.1064	0.131

Column 1: Number of district, ordered with respect to increasing population size; 2: Name of district; 3: FDH efficiency score; 4: Number of dominating districts (including the own one); 5: DEA efficiency score; 6: Mean of bootstrapped DEA efficiency scores; 7: bias-corrected DEA efficiency scores; 8: estimated bias of DEA efficiency scores using the bootstrap; 9: standard deviation of bootstrap DEA efficiency scores; 10: lower bound of 95% bootstrap confidence interval; 11: upper bound of 95% bootstrap confidence interval

From Table 2.1 we can see that the initial DEA efficiency estimators of all districts are well included in the unit interval. Furthermore, only three districts are efficient: Ras Asfour (2), Labkhata (9) and Ain Lehjer (54), and one is close to the frontier with a score equal to 0.9899.

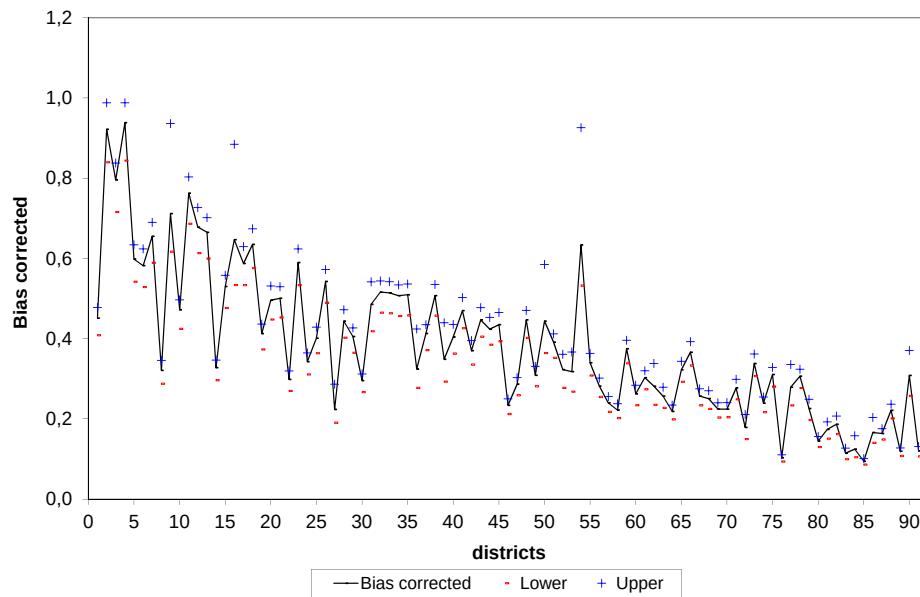
The districts in Table 2.1 are ranked with ascending population size. This suggests that rural districts having a large population size tend to have weak efficiency score and are consequently far from the efficiency frontier as shown in Figure 2.3. For instance, the estimator of the efficiency score has reached 0.1024 for the district of Arekmane (ordered 85th among 91 according to population size). The negative relation between population size and efficiency is confirmed by a Kendall test which strongly rejects that the correlation between the population size and the DEA efficiency estimates is equal to zero with a $p - value < 2.2e - 16$ and indicates an inverse relation between them with a Kendall's tau estimate given by $\tau = -0.6395$. The truncated regression estimation on the Shephard efficiency scores as an environmental variable of the two-stage procedure described in Simar and Wilson (2007) confirm indirectly this relationship.

We used the bias correction given by expression (2.3.3) with $B = 2000$. Denoted by $\tilde{\theta}$, the bias corrected estimators indicate that there are no efficient rural districts and that only 6 percent of the districts are close to the efficiency frontier with a score estimated above 0.70. In addition, the district of Ain Lehjer (54), which was declared efficient by the initial estimator of the DEA efficiency score, is inefficient with an estimator $\tilde{\theta}$ equal to 0.6336, as the bias was important. This indicates that even with a financial autonomy of 1656 percent, this district clearly fails to reach the efficiency frontier. The district of Arekmane (85) recorded the lowest score of efficiency estimated at 0.0944.

In order to test if the bias can be disregarded, the ratio of the estimated bias to the standard error of the bootstrap estimates defined by $\left| \widehat{bias}^*(\hat{\theta}(x, y)) \right| / \hat{\sigma}^*$ has been analyzed. Following a suggestion by Efron (1982), the bias is significant since the ratio exceeds 0.25. It is possible to conclude in our case that the bias correction should be used and the bootstrap bias correction provides more accurate results. On the other hand, we note that the bias corrected estimator for each district is obviously in the corresponding confidence interval, but the

range of the interval is more important for some districts such as Ras Asfour (2), Labkhata (9), Laaouinate (16) and Ain Lehjer (54), which have relatively high bias corrected estimators. This can also be viewed in Figure 2.3.

Figure 2.3: Graph of bias corrected estimators and their confidence interval (Districts ranked according to increasing population size)



It should be noted also that the initial DEA estimators of the efficiency scores are often outside the corresponding confidence interval. They are also close to the upper bound of the interval, since they are positively biased.

FDH results

As pointed out in Table 2.1, eight districts are declared efficient using the FDH approach. This represents only 8.79 percent of the total of the population under study. Nevertheless, this percentage is almost three times higher than that given by the DEA approach. Note also that districts which are declared efficient by DEA approach are also efficient by the FDH approach, which follows

by construction because the FDH approach is less restrictive than that of DEA. On the other hand, the most populated districts have generally a weak FDH efficiency reflecting the same finding as for the DEA approach. The correlation between estimated DEA and FDH efficiency scores is 0.845, which shows that both approaches tend to give similar results where mainly the distance from the frontier differs.

Column 4 of Table 2.1 reports the number of districts which dominate a given district (including the own district). Each FDH inefficient district is dominated by at least another district. For instance, the inefficient district Arekmane (85) is dominated by 46 districts. This means that with the same quantity of resources, ratios of the financial autonomy of the 46 districts exceed that of Arekmane. Furthermore, it can reach efficiency if it reduces its resources by 90 percent, meaning that it can be efficient with only 10 percent of its resources.

2.6 Conclusion

The technical efficiency determination in the input orientation of the Data Envelopment Analysis requires testing returns to scale in order to define the primal linear programming problem. A procedure for the determination of the dual from the primal model was developed for the case of constant returns to scale. Since the efficiency scores are often over-estimated, a bootstrap procedure is used to correct the bias by the mean or by the median. The bootstrap efficiency scores allow us to make statistical inference on the DEA efficiency by using them to build confidence intervals.

In this study, we estimate efficiency scores of the financial autonomy of the Moroccan rural districts in the oriental region for the budgetary year 1998/1999. The inputs are expressed as the operating receipts for the DMU_i to produce the output expressed as the financial autonomy for the same DMU_i . Statistical tests suggested variable returns to scale (VRS) for the data. Bias corrected results indicated that they are well in the unit interval and in the corresponding confidence interval. They indicated also that there are no efficient rural districts and that only 6 percent of the districts are close to the efficiency frontier with a score estimated above 0.70. In addition, the most efficient district is Ain

Lehjer, but even with a financial autonomy of 1656 percent it fails to reach the efficiency frontier.

Being less restrictive than the DEA approach, the FDH analysis delivered efficiency scores generally larger than those of DEA but the DMU ranking is very similar for both approaches.

Finally, we found that generally rural districts having a large population size have a weak efficiency score. If data become available, a detailed analysis using the two-stage procedure described in Simar and Wilson (2007) could be done on the socio-economic and demographic factors such as the geographical distance from the center and the training level of the local council members, which may explain these inefficiencies.

Chapter 3

Inference in stochastic frontier analysis with dependent error terms

3.1 Introduction

Efficiency analysis has often been carried out using nonparametric frontier models such as the Data Envelopment Analysis (DEA) or the Free Disposal Hull (FDH). An alternative approach is to use Stochastic Frontier Analysis (SFA), which includes an error term such that deviations from the frontier can be purely random without necessarily indicating inefficiency. SFA can be formulated both in a parametric or nonparametric framework, but the parametric SFA has certainly been predominant in the literature and in applications. The basic idea of all approaches is the comparison between the Decision Making Units (DMU, firms for example) in order to know how inputs are used to produce outputs and the comparison is based on the Technical Efficiency (TE) score achieved by each unit. By definition, technical efficiency reflects the ability of the firm

to obtain maximal output from a given set of inputs.

The nonparametric frontier approach using DEA or FDH requires minimal assumptions regarding the structure of the production and does not impose restrictions on the functional form relating inputs and outputs. It does not account for noise in the data, so it implicitly assumes that every deviation from the frontier is considered as inefficiency.

However, in the parametric SFA, assumptions have to be made both about the functional form and the distribution of the two types of error, namely, the symmetric stochastic error term and the divergence of observations from the efficient frontier. This stochastic frontier approach in the efficiency analysis was simultaneously and independently introduced by Aigner, Lovell and Schmidt (1977) and by Meeusen and Van den Broeck (1977). Later, several extensions have been proposed by, for example, Agahi, Zarafshani and Behjat (2008), Greene (1980a, 2010), Kumbhakar and Knox Lovell (2000), Simar and Wilson (2010) and Smith (2008). A *FRONTIER* software was developed by Coelli (1996) in order to estimate the stochastic frontier production and the cost function in the case where the two components of the error term are independent. This software is now also available in the statistical computation environment R, see Appendix D and RTeam (2011). As a consequence of its increasing computational availability, stochastic frontier analysis has been widely applied in several areas.

Recently, Smith (2008) has proposed an SFA model allowing for dependence between the two error components. The dependence can be explicitly modelled using copula functions, while maintaining typical assumptions about the marginal distribution of the error terms. Estimation of the model using the Corrected Ordinary Least Squares (COLS) and Maximum Likelihood (ML) methods is straightforward but can be computationally challenging. Furthermore, inference about the technical efficiencies is not standard. In this work, we propose a bootstrap procedure, which is an extension of an algorithm proposed by Simar and Wilson (2010) to the copula case. This allows to obtain not only point estimates, but also confidence intervals for the estimated technical efficiencies.

We apply the model to the estimation of technical efficiencies of Moroccan

municipalities, defining operating receipts as input and financial autonomy as output. The model is estimated with alternative distributions for the one-sided error term, as well as alternative copulas. The best model is selected using classical information criteria. The obtained bootstrap confidence intervals for the technical efficiency estimates are narrow, supporting the adequacy of our methodology and the interpretation of the results. We find that, contrary to common understanding, no municipality in the central regions of the country is close to the frontier.

The remainder of the chapter is organized as follows. The second section gives an overview of parametric SFA and its history, presents the model with dependent error terms and explains the estimation and inference using the bootstrap. The third section presents the application of the proposed methodology, and finally the conclusions in Section 4 will summarize the analysis.

3.2 Parametric stochastic frontier models

Classical parametric stochastic frontier models assume that there is a production function f that converts $X \in \mathbb{R}_+^p$, a vector of inputs of dimension p , into a scalar output $Y \in \mathbb{R}_+$. Supposing that one has n observations of (X_i, Y_i) the model can be written for the i -th DMU as

$$y_i = f(x_i, \beta) + \varepsilon_i, \quad i = 1, \dots, n \quad (3.2.1)$$

where $y_i = \log(Y_i)$, $x_i = \log(X_i)$, β is a vector of parameters of dimension $l + 1$ to be estimated, and ε_i is a stochastic error term. The function $f(x_i, \beta)$ is interpreted as the production frontier.

The stochastic term ε_i contains information about both the noise and the inefficiency. It can be decomposed into a technical inefficiency and a noise term, which can be estimated. In particular, a typical specification is given by

$$\varepsilon_i = v_i - u_i, \quad (3.2.2)$$

where v is a Gaussian error term, $(v_i \sim N(0, \sigma_v^2))$, and u is a stochastic error term with non-negative support ($u_i \geq 0$).

Note that the stochastic component v_i that describes random noise affecting the production process is not directly attributable to the producer or the underlying technology. The noise may come from weather changes, economic adversities, etc. The other component, u_i , measures technical inefficiency in the sense that it measures the shortfall of output y_i from its maximal possible value given by the stochastic frontier ($f(x_i, \beta) + v_i$) and it is equal to zero for a technically efficient decision unit. Then, the one-sided error term $u_i \geq 0$ allows the distinction between DMU (e.g. firms) that are on the frontier ($u_i = 0$) and others that are below the frontier ($u_i > 0$).

The stochastic model then permits to estimate β and its standard errors and, consequently, to make statistical tests of hypotheses. However, one of the criticisms of this model is that there is no *a priori* justification for the selection of the distributional form for u_i . Several choices have been made in the literature, see e.g. the overview of Kumbhakar and Knox Lovell (2000), for example, the exponential, the half-normal, the truncated normal or the Gamma distribution. Furthermore, in order to decompose the error term ε into its two components, one has to make assumptions on their dependence. Classical SFA assumes that they are independent. Let us first recall this approach, see e.g. Jondrow et al. (1982).

3.2.1 Classical SFA with independence

The parameters of the model described by (3.2.1) and (3.2.2) can be estimated using, for instance, the maximum likelihood method and ε_i can be predicted by $\hat{\varepsilon}_i = y_i - f(x_i, \hat{\beta})$, which contains information on u_i . Jondrow et al. (1982) propose a decomposition by considering the expected value of u , conditional on $\varepsilon = v - u$. They proceed by considering the conditional distribution of u_i given ε_i . Either the mean or the mode of this distribution can be used as a predictor of u_i .

In the normal-half-normal case, $v_i \sim N(0, \sigma_V^2)$, u has a half-normal distribution ($u_i \sim N^+(0, \sigma_U^2)$), and v and u are supposed to be independent. Based on these assumptions, one can derive analytical expressions for the marginal distribution of ε and the conditional distribution of u given ε , see Jondrow et al. (1982).

3.2.2 SFA with dependent error components

Consider in this section the simplest cross-section case with n independent DMUs. The most general way to introduce dependence is to use copula functions, which in the present context has been proposed recently by Smith (2008). Appendix B gives a definition and some properties of copulas, and provides frequently used examples of parametric copula functions.

Let us consider the normal, half-normal production frontier model with Cobb-Douglas production function. Thus, let $v \sim i.i.d.N(0, \sigma_V^2)$ and $u \geq 0$, $u \sim i.i.d.N^+(0, \sigma_U^2)$ with $E(u) = \sigma_U \sqrt{2/\pi}$ and $Var(u) = ((\pi - 2)/\pi) \sigma_U^2$. We know also that $\varepsilon = v - u$, and hence $Var(\varepsilon) = Var(u) + Var(v) - 2Cov(u, v)$. Therefore, a positive correlation between u and v (i.e. $Cov(u, v) > 0$) reduces the variance of ε , and a negative correlation increases it.

The joint density of u and v when they are dependent is expressed for all $u \geq 0$ and $v \in \mathbb{R}^n$ by

$$g_\theta(u, v) = f_1(u)f_2(v)c_\theta(F_1(u), F_2(v))$$

In the following we give two examples. Consider first the Gaussian copula (see Appendix B.1), for which the joint density of u and v can be derived as

$$g_\theta(u, v) = \frac{1}{\pi\sigma_U\sigma_V} \exp\left\{-\frac{1}{2\sigma_U^2}u^2 - \frac{1}{2\sigma_V^2}v^2\right\} \times \left[\frac{\phi_{2,\theta}(\Phi^{-1}(F_1(u)), \Phi^{-1}(F_2(v)))}{\phi(\Phi^{-1}(F_1(u))) \cdot \phi(\Phi^{-1}(F_2(v)))}\right],$$

where $\theta \in [-1, 1]$ is the parameter of the copula, $F_1(u) = 2\Phi(u/\sigma_U)$ and $F_2(v) = \Phi(v/\sigma_V)$.¹

For the case of an FGM copula (see Appendix B.2), the joint density be-

¹Note that

$$F_1(u) = \int_{-\infty}^u \frac{2}{\sigma_U \sqrt{2\pi}} \exp\left\{-\frac{1}{2}\left(\frac{t}{\sigma_U}\right)^2\right\} dt = \int_{-\infty}^{u/\sigma_U} \frac{2}{\sqrt{2\pi}} \exp\left\{-\frac{1}{2}z^2\right\} dz = 2\Phi\left(\frac{u}{\sigma_U}\right).$$

comes

$$g_{\theta}(u, v) = \frac{1}{\pi\sigma_U\sigma_V} \exp \left\{ -\frac{1}{2\sigma_U^2}u^2 - \frac{1}{2\sigma_V^2}v^2 \right\} \\ \times \left[1 + \theta - 4\theta\Phi\left(\frac{u}{\sigma_U}\right) - 2\theta\Phi\left(\frac{v}{\sigma_V}\right) + 8\theta\Phi\left(\frac{u}{\sigma_U}\right)\Phi\left(\frac{v}{\sigma_V}\right) \right],$$

where $\theta \in [-1, 1]$. The joint density of u and ε is obtained by replacing v in $g_{\theta}(u, v)$ by $v = \varepsilon + u$.

To obtain the density of ε , the joint density of u and ε is then integrated by the variable u ,

$$\begin{aligned} g_{\theta}(\varepsilon) &= \int_0^{+\infty} g_{\theta}(u, \varepsilon) \, du \\ &= \int_0^{+\infty} f_1(u)f_2(\varepsilon + u)c_{\theta}(F_1(u), F_2(\varepsilon + u)) \, du \\ &= E_U(f_2(\varepsilon + u)c_{\theta}(F_1(u), F_2(\varepsilon + u))) \end{aligned} \quad (3.2.3)$$

If no analytical solution of the integral is available, one can approximate it numerically via simulation by drawing a large number m of random variables U from the marginal distribution of u , and calculate for any value ε ,

$$g_{\theta}(\varepsilon_i) \cong \frac{1}{m} \sum_{j=1}^m f_2(\varepsilon_i + u_j)c_{\theta}(F_1(u_j), F_2(\varepsilon_i + u_j))$$

Replacing ε by $y - f(x, \beta)$ in the expression of $g_{\theta}(\varepsilon)$ gives the density of y . Assuming independence across DMUs, the log-likelihood function is given by

$$l(\vartheta) = \sum_{i=1}^n \log g_{\theta}(\varepsilon_i) = \sum_{i=1}^n \log g_{\theta}(y_i - f(x_i, \beta)), \quad (3.2.4)$$

where $\vartheta = (\sigma_U, \sigma_V, \theta, \beta)'$, and the ML estimator of ϑ is defined as

$$\hat{\vartheta}_{ML} = \arg \max_{\vartheta \in \Theta} l(\vartheta)$$

Maximization of the log-likelihood function is typically done using numerical techniques, as analytical solutions are rarely available.

Based on ML parameter estimates, one can address the issue of estimating technical efficiencies, defined as the expectation of efficiency conditional on observed residuals, see Battese and Coelli (1988). Hence, technical efficiency of DMUs, which depends on the parameter $\vartheta = (\sigma_U, \sigma_V, \theta, \beta)$ and on the observed input x and output y , is defined by

$$TE_{\vartheta} = E\left(\exp\{-U\} \mid \varepsilon\right)$$

Using the marginal distribution of ε in (3.2.3), and the joint density of u and ε we can calculate

$$\begin{aligned} TE_{\vartheta} &= \int_{\mathbb{R}^+} \exp\{-u\} g_{\theta}(u \mid \varepsilon) du \\ &= \frac{1}{g_{\theta}(\varepsilon)} \int_{\mathbb{R}^+} \exp\{-u\} g_{\theta}(u, \varepsilon) du \\ &= \frac{E_U\left(\exp\{-u\} \cdot f_2(u + \varepsilon) \cdot c_{\theta}\left(F_1(u), F_2(u + \varepsilon)\right)\right)}{E_U\left(f_2(u + \varepsilon) \cdot c_{\theta}\left(F_1(u), F_2(u + \varepsilon)\right)\right)} \end{aligned} \quad (3.2.5)$$

For a given ϑ , this expression can again be approximated via simulation by drawing a large number of random variables U and approximate the expectations appearing in numerator and denominator by the corresponding simulation means. Replacing ϑ by its ML estimator provides the ML estimator of TE_{ϑ} .

While point estimation of TE_{ϑ} is straightforward, although it may be computationally demanding, it is less obvious how to obtain interval estimates and how to do inference. We next describe an algorithm for obtaining confidence intervals for the technical efficiencies.

3.2.3 Bootstrap confidence intervals for technical efficiencies

We now propose a statistical inference procedure to construct confidence intervals for technical efficiencies in the SFA model with dependence. We use an extension of the bootstrap procedure described in Simar and Wilson (2010). In particular, step 2 of algorithm#3 of Simar and Wilson (2010) is modified to

take into account the dependence between u_i^* and v_i^* using the Clayton copula. The various steps of the algorithm are as follows:

Step 1. Estimate $\vartheta = (\sigma_U, \sigma_V, \theta, \beta_0, \beta_1)$ according to (3.2.4), using (x_i, y_i) , $i = 1, \dots, n$ and using numerical optimization procedures to get $\hat{\vartheta} = (\hat{\sigma}_U, \hat{\sigma}_V, \hat{\theta}, \hat{\beta}_0, \hat{\beta}_1)$ and to compute $TE_{\hat{\vartheta}}$.

Step 2. For $i = 1, \dots, n$, draw $u_i^* \sim N^+(0, \hat{\sigma}_U^2)$ and $v_i^* \sim N(0, \hat{\sigma}_V^2)$ such that their dependence is given by the Clayton copula, and then compute $y_i^* = f(x_i, \hat{\beta}) + v_i^* - u_i^*$.

There are several procedures to generate the pair (u_i^*, v_i^*) with dependence given by the Clayton copula, we mention one of them which uses the conditional distribution approach described in Nelsen (1999), page 41 and denoted $c_{w_1}(w_2)$,

$$\begin{aligned} c_{w_1}(w_2) &= P(W_2 \leq w_2 \mid W_1 = w_1) \\ &= \lim_{\Delta w_1 \rightarrow 0} \frac{C(w_1 + \Delta w_1, w_2) - C(w_1, w_2)}{\Delta w_1} \\ &= \frac{\partial C(w_1, w_2)}{\partial w_1}. \end{aligned} \quad (3.2.6)$$

The four steps of this procedure are:

- (a) Draw two independent uniform random variables (w_{1i}, t_{2i}) such that $w_{1i} \sim U(0, 1)$ and $t_{2i} \sim U(0, 1)$.
- (b) Set $w_{2i} = \left[w_{1i}^{-\hat{\theta}} \left(t_{2i}^{-\hat{\theta}/(1+\hat{\theta})} - 1 \right) + 1 \right]^{-1/\hat{\theta}}$.
- (c) Set $u_i^* = F_1^{-1}(w_{1i})$ and $v_i^* = F_2^{-1}(w_{2i})$, where F_1 and F_2 are the cumulative distribution function of the $N^+(0, \hat{\sigma}_U^2)$ and $N(0, \hat{\sigma}_V^2)$ respectively.
- (d) Repeat steps a) to c) to generate n pairs (u_i^*, v_i^*) .

Step 3. Using the pseudo-data $\mathcal{S}_{b,n}^* = \{(x_i, y_i^*)\}_{i=1}^n$, compute a bootstrap estimate $\hat{\vartheta}_b^* = \arg \max_{\vartheta \in \Theta} l(\vartheta \mid \mathcal{S}_{b,n}^*)$ after replacing y_i by y_i^* in (3.2.4) and then compute a bootstrap estimate $\widehat{TE}_b^*$ using (3.2.5) after replacing ε by $\varepsilon_b^* = y - f(x, \hat{\beta}_b^*)$, where x and y represent the observed data.

Step 4. Repeat steps 2 and 3 B times to obtain estimates $\mathcal{B}^* = \{\hat{\nu}_b^*\}_{b=1}^B$. Then, use $\mathcal{B}^*$ to obtain the set of B bootstrap estimates of technical efficiency, $\mathcal{E}^* = \{\widehat{TE}_b^*\}_{b=1}^B$. For each individual i (row i of the $\mathcal{E}^*$ matrix, denoted by $\mathcal{E}_i^*$), $i = 1, \dots, n$, compute the $(\frac{\alpha}{2})$ and the $(1 - \frac{\alpha}{2})$ quantiles for $\mathcal{E}_i^*$ by considering its B components. The $100 \times (1 - \alpha)$ percentile bootstrap confidence interval of the statistic of interest TE is obtained by the probability $P\left((\mathcal{E}_i^*)_{\frac{\alpha}{2}} < TE_i < (\mathcal{E}_i^*)_{1-\frac{\alpha}{2}}\right) = 1 - \alpha$. Hence, using the $100 \times (\frac{\alpha}{2})$ and $100 \times (1 - \frac{\alpha}{2})$ percentiles, we define the lower and the upper bounds of the confidence interval as $TE_i \in \left[(\mathcal{E}_i^*)_{\frac{\alpha}{2}}, (\mathcal{E}_i^*)_{1-\frac{\alpha}{2}}\right]$.

The proposed algorithm may be computationally intensive but it is straightforward to apply and to implement. Furthermore, bootstrap techniques have the advantage of taking implicitly the estimation uncertainty of the parameters into account.

We investigate the performance of the proposed method in a Monte Carlo study. We use the same model and parameters as in Simar and Wilson (2010), i.e. $\beta_0 = \log(10)$, $\beta_1 = 0.8$, $\lambda^2 = \sigma_U^2/\sigma_V^2 = 2$, and quantiles of ε given by $\varepsilon_{(q)}$, where $q \in \mathcal{Q}$ and $\mathcal{Q} = \{0.1, 0.3, 0.5, 0.7, 0.9\}$. The input value is fixed at $x_0 = 60$, and five log-output values are given by $y_{(q)} = \exp(\beta_1 + \beta_2 x_0 + \varepsilon_{(q)})$, corresponding to the five quantiles of $\varepsilon_{(q)}$. The dependence parameter θ of the Clayton copula is fixed at 2 which means that the association between the two error terms is evaluated at 0.5 according to the Kendall's τ . We use $n = 100$ and $n = 1000$ as sample sizes, $M = 1000$ Monte Carlo trials, and $B = 500$ bootstrap replications. The estimated coverages of the bootstrap confidence intervals are given in Table 3.1. The results for the independence case ($\theta = 0$) closely match those of Simar and Wilson (2010), as expected. A general result is that, for the given model, coverage is slightly higher under positive dependence of U and V than under independence. There is over-coverage especially for lower quantiles and $n = 100$, which almost disappears both under dependence and independence for $n = 1000$. On the other hand, there is under-coverage for the upper quantile 0.9 and $n = 100$, which is less severe under dependence, and which also disappears for $n = 1000$. Thus, the conclusion is that in large samples, the proposed bootstrap algorithm is a reliable method to construct confidence intervals under dependence, while in small samples, the over- or

Table 3.1: Estimated coverages of $E[e^{-U}|\theta, x_0 = 60, y_0 = y_{(q)}]$ by bootstrap confidence intervals.

$\theta = 0$ (independence)						
$n = 100$			$n = 1000$			
$1 - \alpha$			$1 - \alpha$			
quantile	0.9	0.95	0.99	0.9	0.95	0.99
0.1	0.931	0.958	0.984	0.892	0.945	0.987
0.3	0.936	0.963	0.982	0.890	0.946	0.985
0.5	0.933	0.964	0.989	0.887	0.943	0.986
0.7	0.884	0.946	0.986	0.879	0.939	0.984
0.9	0.704	0.794	0.918	0.870	0.928	0.975
$\theta = 2$						
$n = 100$			$n = 1000$			
$1 - \alpha$			$1 - \alpha$			
quantile	0.9	0.95	0.99	0.9	0.95	0.99
0.1	0.952	0.980	0.991	0.910	0.949	0.989
0.3	0.954	0.983	0.992	0.917	0.956	0.996
0.5	0.949	0.975	0.990	0.908	0.953	0.992
0.7	0.912	0.960	0.985	0.900	0.948	0.990
0.9	0.805	0.894	0.937	0.891	0.942	0.983

under-coverage of the bootstrap may depend on the considered quantile of the output variable.

3.3 An illustrative analysis of the efficiency of Moroccan municipalities

In this illustrative analysis, we study the management of the financial resources of Moroccan municipalities, which in times of tight public budgets is an important issue for policy makers. We consider 1298 Moroccan rural municipalities (DMUs) to produce one output which is the financial autonomy using operating receipts as input for the budgetary year 1998/1999. The operating receipts include ten receipts which are the urban tax, the tax on the collection of waste, the tax of the licence, the forest domain product, the taxes and assimilated taxes, the product of services, the product and the income of goods, the concessions, the subsidies and competition and finally the order receipts. We chose to use the aggregate measure of operating receipts as the single input because it is

a meaningful variable and because it allows us to reduce the dimensional complexity. Moreover, the operating receipts variable is highly correlated with five of the ten inputs. Some preliminary analysis with a multiple input framework did not change the main conclusions drawn from the aggregate input case.

Furthermore, financial autonomy is defined as the ratio of the own receipts and the operating expenses. As for the own receipts, they include all operating receipts except the subsidies and competition. After the decentralization of the Moroccan administration, this kind of data is not available after 1999. Thus, our data consists of pairs (X_i, Y_i) where X_i represents the single input expressed by the operating receipts of the DMU_i used to produce the output Y_i , i.e. the financial autonomy of the same DMU_i .

Municipalities are clustered in provinces and regions, and we could have added a hierarchical structure to the model, but did not pursue this direction for simplicity. In the interpretation of the results, we will come back to this point and try to interpret the estimated DMU efficiencies with respect to their geographical and political situation.

To represent the production technology, we consider the frequently used Cobb-Douglas and translog production functions, see Section 1.1 or e.g. Christensen, Jorgenson and Lau (1971) for a general definition of the translog production function. The Cobb-Douglas function is nested in the translog one, such that it can be tested. In our case with just one input variable, the test reduces to testing a linear model against a quadratic one.

Our distributional assumptions about the error terms are as follows. For the random noise term v we assume a normal distribution, while either a half-normal (HN) or truncated normal (TN) distribution are adopted for the inefficiency component u . Again, the HN model is nested in the TN model, such that it can be tested easily. The general model with translog production function and TN distribution for u then reads

$$y_i = \beta_0 + \beta_1 x_i + \beta_2 x_i^2 + v_i - u_i, \quad i = 1, \dots, n \quad (3.3.1)$$

where

$$v_i \sim N(0, \sigma_V^2), \quad u_i \sim N^+(\mu, \sigma_U^2).$$

Table 3.2: Technical Efficiency and Log-Likelihood values in the independence case

Name(Pop. Order)	CD-HN	TL-HN	CD-TN	TL-TN
Tagante(244)	0.919	0.917	0.916	0.915
Tidili Mesfioua(1239)	0.881	0.880	0.887	0.888
Agafay(667)	0.848	0.850	0.861	0.863
Timzguida-Ouftas(185)	0.831	0.834	0.848	0.850
Adaghas(102)	0.819	0.819	0.839	0.840
Bouhmame(1252)	0.814	0.817	0.835	0.837
Ida ou Guelloul(313)	0.809	0.811	0.831	0.834
Taghazout(234)	0.809	0.811	0.830	0.833
Mouarid(299)	0.807	0.808	0.830	0.832
Ait Aissi Ihahane(257)	0.805	0.807	0.828	0.830
⋮	⋮	⋮	⋮	⋮
Jdiriya(46)	0.021	0.021	0.021	0.021
Tifariti(69)	0.021	0.020	0.020	0.021
Haouza(77)	0.021	0.020	0.020	0.020
log-likelihood	-1771.50	-1771.165	-1764.131	-1764.117
mean	0.478	0.475	0.558	0.560

HN: half-normal; TN: truncated normal; CD: Cobb Douglas; TL: translog

The special case Cobb-Douglas is attained by restricting $\beta_2 = 0$, and HN by restricting $\mu = 0$. The stochastic frontier model in (3.3.1) with independent u and v is estimated by Maximum Likelihood, where initial values are set to the Ordinary Least Squares (OLS) estimates. The OLS estimators of β_1 and β_2 are unbiased, and the intercept can be bias-adjusted using the Corrected Ordinary Least Square (COLS) method, see e.g. Greene (1980b).

Assuming independence of u and v , the results summarized in Table 3.2 reveal that models with the translog function have slightly higher log-Likelihood values compared with the corresponding Cobb-Douglas models. Using a likelihood ratio test statistic, the null hypothesis $\beta_2 = 0$ can not be rejected with a p-value equal to 0.4130 in the half-normal case, and 0.6079 in the truncated normal case. Therefore, we pursue our analysis accepting the Cobb-Douglas specification. In addition, testing $\mu = 0$ in the truncated normal model does not lead to a rejection at a level of 5%, which means that we accept the half-normal specification. Thus, the Normal-HN model with Cobb-Douglas production function is considered and its estimates will be chosen as initial values in the case of dependence between u and v . They are reported in Table 3.3.

As it is not possible to directly estimate the error components u and v , but

Table 3.3: Maximum likelihood estimates in the independence case for the CD-HN model

	Estimate	Std. Error	t value	p.value
β_0	-12.0559	0.6036	-19.975	< 2.2e-16
β_1	1.1149	0.0410	27.196	< 2.2e-16
σ^2	1.8322	0.1184	15.474	< 2.2e-16
γ	0.7796	0.0321	24.323	< 2.2e-16
λ	1.8808			
σ_U	1.1952			
σ_V	0.6355			

$\lambda, \sigma_U, \sigma_V$ computed from $\lambda = \frac{\sigma_U}{\sigma_V}$, $\sigma^2 = \sigma_U^2 + \sigma_V^2$ and $\gamma = \frac{\lambda^2}{1+\lambda^2}$

only their difference, $v - u$, we can not directly test the independence between them. However, it is possible to generalize the preferred model under independence, i.e. CD-HN, to allow for dependence. In particular, introducing a copula which nests the product copula, i.e. independence, as a special case allows to test the null hypothesis of independence using a likelihood ratio test, or use a model selection criterion such as AIC or BIC to distinguish between the two models. In our case, it will turn out that copula models significantly outperform the model assuming independence, which indicates that the assumption of independence is too restrictive. Ignoring dependence between the error components may lead to biased estimates of β , σ_U and σ_V . In the following, we therefore consider various copula models for the joint distribution of u and v .

The maximization of the log-likelihood function in (3.2.4) often requires numerical derivatives. We use the *mle* function from the **stats4** package in the R software. The estimates under independence reported in the Table 3.3 are used as initial values. Several optimization methods as a variant of a simulated annealing method (SANN) given in B  lisle (1992) and the Nelder-Mead method given in Nelder and Mead (1965) are offered in the R package, but the selected one is the Nelder-Mead method.

As pointed out by Ritter and Simar (1997) and Simar and Wilson (2010), the estimation with location parameter μ in the TN model can be numerically difficult and may require very large sample sizes. The reason is that, for moderate sample sizes, the likelihood function is flat with respect to μ , such that a practical identification issue arises, although asymptotically the model

is well identified. An alternative is to set this parameter to a predetermined value. In our case, we have experimented with several values and have chosen to set $\mu = -1$ for the models with truncated normal distribution, as it gave in most cases the best fit. Technical efficiencies are estimated according to (3.2.5) for ten models using the Cobb-Douglas function, the normal distribution for the noise term v , the half-normal and the truncated normal distributions with $\mu = -1$ for the inefficiency error u and using five copulas. These copulas are the Ali-Mikhail-Haq (AMH), Clayton, Fairlie-Gumbel-Morgenstern (FGM), Frank and Gaussian copulas.

Table 3.4: Technical Efficiency and Log-Likelihood values with dependence

N	HN-AMH	HN-FGM	HN-Clay	HN-Frank	HN-Gauss	TN-AMH	TN-FGM	TN-Clay	TN-Frank	TN-Gauss
244	0.835	0.801	0.748	0.903	0.788	0.838	0.849	0.779	0.874	0.742
1239	0.751	0.722	0.647	0.850	0.752	0.758	0.779	0.695	0.813	0.716
667	0.686	0.674	0.577	0.803	0.723	0.695	0.727	0.633	0.762	0.695
185	0.658	0.660	0.550	0.779	0.709	0.667	0.708	0.607	0.739	0.684
102	0.640	0.654	0.533	0.763	0.701	0.651	0.701	0.594	0.726	0.675
1252	0.633	0.653	0.527	0.757	0.695	0.642	0.697	0.586	0.718	0.673
313	0.627	0.651	0.522	0.751	0.692	0.636	0.694	0.581	0.714	0.669
234	0.625	0.651	0.520	0.749	0.690	0.633	0.693	0.578	0.711	0.669
299	0.623	0.651	0.519	0.748	0.690	0.633	0.694	0.579	0.712	0.668
...	...	...	...	...	...	...	...	...	...	...
46	0.006	0.014	0.004	0.012	0.011	0.007	0.008	0.005	0.018	0.005
77	0.006	0.014	0.004	0.012	0.011	0.007	0.008	0.005	0.018	0.005
69	0.006	0.013	0.004	0.011	0.011	0.007	0.008	0.005	0.017	0.004
Mean	0.382	0.426	0.342	0.421	0.423	0.459	0.454	0.394	0.471	0.407
Median	0.452	0.462	0.435	0.452	0.460	0.535	0.500	0.487	0.516	0.455
θ	0.965	0.984	1.375	1.618	0.348	0.981	0.960	1.259	1.989	0.501
λ	2.329	1.920	2.528	2.276	1.986	2.619	2.574	2.649	2.130	2.694
logLik	-1752.811	-1761.141	-1750.948	-1759.379	-1767.649	-1755.447	-1758.865	-1756.608	-1759.623	-1763.715

 $\mu = -1$ for the truncated normal distribution

Table 3.4 reports a subsample of the estimated technical efficiencies for the ten alternative models, together with the estimated θ of the corresponding copula. Furthermore, the ratio of standard deviations $\lambda = \sigma_U/\sigma_V$ is reported, which is a measure for the asymmetry of the composite error distribution. Finally, the estimated likelihood values are reported. Since all models have the same number of parameters, standard model selection criteria such as AIC or BIC are equivalent to choosing the model with the highest likelihood value. In our case, this is the model where the error term v has a Normal distribution, the inefficiency term u has a half-normal distribution and where the dependence between u and v is expressed by the Clayton copula. This preferred model will be denoted HN-Clay.

In order to validate the chosen model, model identification based on maximum likelihood is considered. Our stochastic model can be written basically as $\log(Y) = \alpha + v - u$, where α is a constant. We know that v is a normal r.v. with constant mean equal to zero ($v \sim N(0, \sigma_V^2)$) and u is a half-normal with constant mean ($u \sim N^+(0, \sigma_U^2)$) where σ_U is the standard deviation parameter of the standard normal. As the distribution of the subtraction of the inefficiency term from the noise term is another, distinct distribution, and the subtraction of their means is fixed and unique, and no a priori knowledge on the noise variance is necessary, the likelihood is uniquely determined and the model is identifiable.

The parameter estimates of the HN-Clay model are presented in Table 3.5. Note that all parameters are significantly different from zero at all usual significance levels. In particular, the independence hypothesis $\theta = 0$ is clearly rejected. To interpret the association between u and v we calculate Kendall's τ which is the probability of concordance minus the probability of discordance, and is thus standardized to the interval $[-1, 1]$. For the Clayton copula, Kendall's τ is given by $\tau = \theta/(\theta + 2)$ and for the estimated $\theta = 1.375$ takes the value $\tau = 0.407$. Clearly, the probability of concordance is higher than the probability of discordance for the random variables u and v . The Appendix B gives Kendall's τ for some copulas.

To classify the estimation results with respect to the 1298 districts, we rank these with respect to increasing population size. Note from the results in Table

Table 3.5: ML estimator of ϑ for the HN-Clay model

ϑ	Estimate	Std. Error	$t_{stat} = \frac{\hat{\vartheta}_k}{SE(\hat{\vartheta}_k)}$
β_0	-11.6612	0.0017	-6859.530
β_1	1.1290	0.0037	305.135
θ	1.3750	0.0027	509.259
σ_U	1.9167	0.0018	1064.833
σ_V	0.7583	0.0017	446.059
logLik	-1750.948		

3.4 that the ranking of efficiency estimates is almost always the same irrespective of the model. No district is close to the frontier, the highest efficiency is attained for the Tagante (244) district, which is in the Guelmim Province (50) in the south of Morocco and which is ranked 244-th according to (increasing) population size. The following districts are Tidili Mesfioua (1239) in the Al Haouz Province (20), Agafay (667) in Agadir Idaoutanane Province (16), Timzguida-Ouftas (185) in Essaouira Province (47), Adaghas (102) in Essaouira Province (47), and Bouhmame (1252) in Safi Province (44). For the chosen HN-Clay model, the Tagante (244) district, for instance, could reach efficiency by reducing its resources by 25.2 percent. The least efficient of all districts is Tifariti (69) in the Es-Smara province(48), with receipts covering less than 1 percent of its expenses.

Note also that there is a very big disparity between districts of the Guelmim-Es Semara region in so far as it includes the most efficient district as well as the three least efficient ones. However, they belong to completely different provinces which are Guelmim for the first one and Es-Semara for the three last ones and both provinces have different geographical specificities. On the other hand, among the ten most efficient districts, seven are in the Marrakech-Tensift-Al Haouz region and among these seven municipalities, five are in the Essaouira Province. Even if they are well classified, their estimates of technical efficiency remain quite far from the frontier irrespective of the model used for the dependence.

It is surprising not to find among the most efficient municipalities those of the central regions (for example the Rabat-Salé-Zemmour-Zaër or the Grand

Casablanca regions) which are close to the central administration and where local council members typically have a high training level. However, in the absence of data on the geographical distance and on the training level of the local elected officials, their effects on the municipality efficiency cannot be formally tested.

Except for the models with Frank and Gaussian copula, efficiency means and medians of all models where the inefficiency term u has a half-normal distribution are lower than those using the truncated normal distribution. Note that the choice of copula affects the estimated efficiency level. The highest levels are obtained for the Frank copula, the lowest for Clayton (HN) and Gaussian (TN). Note also that the medians are higher than the means in both cases, reflecting the fact that both distributions are negatively skewed in all cases. This can be seen in Figure 3.1, which displays the Box-plots of estimated efficiencies for all models and which illustrate the dispersion and skewness of their distributions.

Figure 3.1: Boxplot of TE for ten models with HN and TN($\mu = -1$)

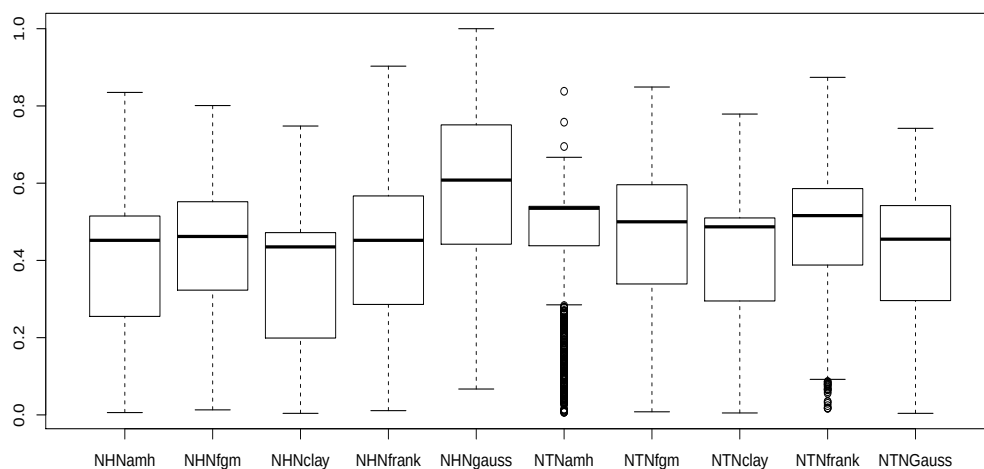
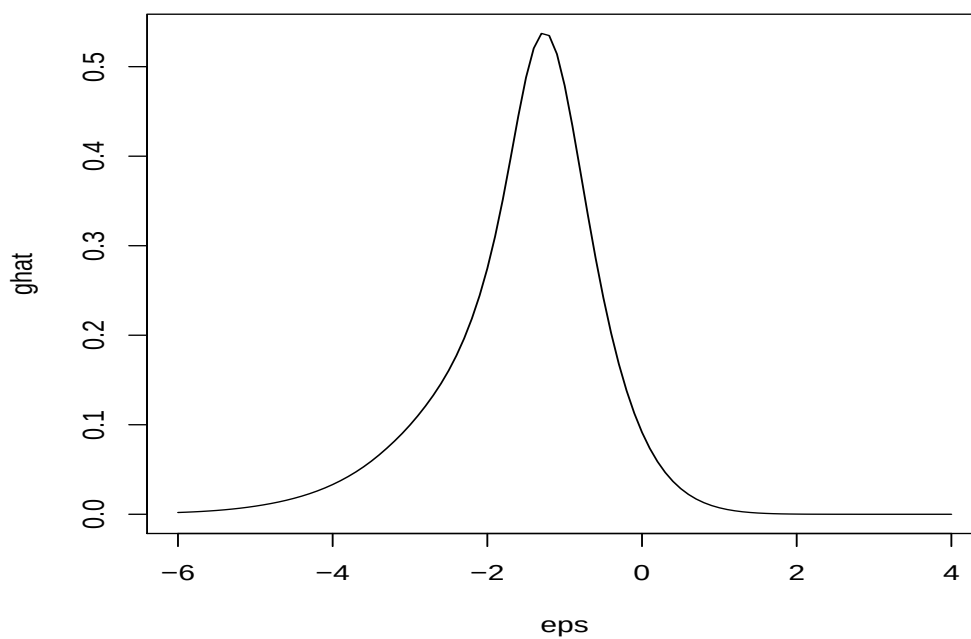


Figure 3.2 depicts the estimated $g_{\theta}(\varepsilon)$ density for the HN-Clay model. It clearly shows the skewness, which can also be expressed in terms of the estimated ratio of standard errors, $\lambda = \sigma_U/\sigma_V$, taking a value 2.528, recalling that

a symmetric density would be obtained for $\lambda = 0$. Note also that due to the dependence between the two error components, the mode of the density is not at zero but shifted to the left. Besides, introducing a dependence between u and v , which is positive in our case, tends to reduce the general level of estimated technical efficiencies, which can be seen by comparing the results of Table 3.4 with those of the independence case in Table 3.2.

Figure 3.2: $g_\theta(\varepsilon)$ distribution for the N-HN-Clay Copula model



We now provide estimated 95% confidence intervals for the technical efficiencies using the bootstrap algorithm presented in Section 3.2.3. Table 3.6 gives an overview of the lower and the upper bounds of bootstrap confidence intervals for the N-HN-Clay model with number of bootstrap replications $B = 700$. As summarized in this table, each estimated efficiency is covered by the associated

Table 3.6: TE confidence intervals for the HN-Clay model

Name(Pop. Order)	Province	Region	Lower	TE	Upper
Tagante(244)	Guelmim	Guelmim-Es Smara	0.6933	0.7484	0.8753
Tidili Mesfioua(1239)	Alhouz	Marrakech-T.-H. ^c	0.5863	0.6468	0.8111
Agafay(667)	Marrakech M. ^a	Marrakech-T.-H. ^c	0.5183	0.5772	0.7558
Timzguida-Ouftas(185)	Essaouira	Marrakech-T.-H. ^c	0.4939	0.5495	0.7284
Adaghas(102)	Essaouira	Marrakech-T.-H. ^c	0.4813	0.5331	0.7145
Bouhmame(1252)	El Jadida	Doukkala-Abda	0.4752	0.5269	0.7039
Ida ou Guelloul(313)	Essaouira	Marrakech-T.-H. ^c	0.4712	0.5217	0.6985
Taghazout(234)	Agadir I. O. ^b	Sous-Massa-Draâ	0.4695	0.5202	0.6945
Mouarid(299)	Essaouira	Marrakech-T.-H. ^c	0.4693	0.5190	0.6959
Ait Aissi Ihahane(257)	Essaouira	Marrakech-T.-H. ^c	0.4674	0.5169	0.6925
...	...	...	...	...	...
Ain Blal(237)	Settat	Chaouia-Ouadigha	0.4388	0.4741	0.6034
Amerzgane(611)	Ouarzazate	Sous-Massa-Draâ	0.4483	0.4739	0.5377
Sidi Lahsen(773)	Taourirt	Oriental	0.4483	0.4739	0.5375
...	...	...	...	...	...
Jdiriya(46)	Es-Semara	Guelmim-Es Smara	0.0032	0.0039	0.0066
Haouza(77)	Es-Semara	Guelmim-Es Smara	0.0032	0.0038	0.0066
Tifariti(69)	Es-Semara	Guelmim-Es Smara	0.0031	0.0037	0.0063

^aMarrakech M.: Marrakech Menara,^bAgadir I. O.: Agadir Ida Outanane,^cMarrakech-T.H.: Marrakech-Tensift-Al Haouz.

confidence interval as expected, and generally the range of each confidence interval is rather small. Note that the estimated efficiencies are generally closer to the lower limit of the interval and, hence, the intervals are not symmetric around the estimated TE values, which is also as expected.

It may be of further interest to discover any links of estimated technical efficiencies with observed characteristics such as the population size. For the selected model HN-Clay we use Kendall's independence test between technical efficiencies and population size, which yields a statistic $\tau = -0.0468$ and corresponding p-value equal to 0.0011. Thus we reject the null hypothesis of independence. The relation between the two variables is opposite, so that highly populated districts tend to be less efficient. This may suggest that population size influences financial autonomy, in which case it could be included in the model. Policymakers may address the problem of increasing the financial autonomy of highly populated districts and those of the central regions.

3.4 Conclusion

In the framework of a stochastic frontier analysis with dependence between the noise term V and inefficiency U , we introduce a bootstrap procedure to

estimate confidence intervals for technical efficiencies. Applying the model to the financing of Moroccan rural districts, we find that estimated technical efficiencies allowing for dependence through copulas tend to be lower than under independence, while the ranking remained basically the same. Furthermore, the most efficient districts are in the regions of Guelmim-ES Semara, Marrakech-Tensift-El Haouz, Sous-Massa-Draâ and Doukkala-Abda and, contrary to prior expectations, no districts of the central regions is among the top classified. We find a significant negative link between estimated technical efficiencies and population size, indicating that highly populated districts tend to be less efficient. Future research may provide a detailed analysis of the socio-economic and demographic factors that could explain inefficiencies such as the geographical distance from the center and the training level of the local councils members.

Chapter 4

Statistical inference for panel stochastic frontier analysis with an illustration of Moroccan drinking water performance

4.1 Introduction

When panel data are available, it is obviously recommended to use the structure of the data to estimate technical efficiencies in the Stochastic Frontier Analysis (SFA) because a panel contains more information than a single cross section. Furthermore, as noted in Schmidt and Sickles (1984) and Kumbhakar and Lovell (2000) some strong distributional assumptions used in the cross-sectional data case can be relaxed with the panel data and the technical efficiency can be estimated consistently when T , the number of time observations for each Decision Making Unit (DMU), is large. Hence, repeated observations can constitute a substitute for some strong distributional assumptions. They can constitute also

a substitute for the independence assumption between the technical inefficiency term and the regressors.

An overview of the research on panel SFA models reveals that Jondrow et al. (1982) generalized the cross-sectional model to the panel data model and used the conditional expectation of the inefficiency term u given the realized value of the error ϵ_i to estimate efficiency. Schmidt and Sickles (1984) and Kumbhakar and Lovell (2000) have proposed, when the panel is balanced¹, models where they supposed that technical efficiency varies across producers but is either constant or varies through time for each producer. Battese, Heshmati and Hjalmarsson (2000) adopted an unbalanced panel to investigate efficiency of labour in the Swedish banking industry. Kim and Lee (2006) assumed a time varying pattern of technical efficiency movements to analyze the productivity growth of several East Asian countries over a period of twenty years. However, all studies handled the panel SFA model with independence between noise and inefficiency terms. Recently, Smith (2008) handled the panel data model with dependent error components using a simulated example but without making inference on the estimated efficiency.

Furthermore, several studies have been done to assess the performance of the water services such as Faria et al. (2005) and Tupper and Resende (2004) which compare the technical efficiency of Brazilian public and private companies in water supply; Sampaio, Barros and Ramajo (2005) which deals with the cost efficiency of the public water service in Portugal, and Vishwakarma and Kulshrestha (2010) which analyses the water supply utility of urban cities in India using the stochastic production frontier analysis. However, most of these studies use cross-sectional data. In the absence of such study in the water domain based on Moroccan data, the performance of the entities responsible of the water management in all regions will be measured by estimating the efficiencies in case of panel data and a nonparametric confidence interval will be proposed.

The objective of our study is to deal with the dependence of the error terms in the panel SFA approach using some proposed time varying models and hence

¹Balanced panel: Each DMU is observed T times.

Unbalanced panel: DMU_i is observed $T_i < T$ times, with not all T_i equal, $i = 1, \dots, n$.

to evaluate the efficiency and to compare the DMUs performances through an empirical data set on the water management in Morocco for the chosen model. So, in this work the production frontiers and panel data are considered to estimate technical efficiency when the two components of the error term are dependent. Efficiency being estimated, statistical inference is needed to draw reliable conclusions. Hence, this work presents also an associated procedure to build confidence intervals on the efficiencies in this considered case. Indeed, the remainder of the chapter is organized as follows: Section 2 describes the model with the copula function, Section 3 presents the procedure of statistical inferences on the Technical Efficiency (TE) measure in the case of panel data with dependent error terms, and the last section presents results of an empirical analysis of the water area in all Moroccan regions with the numerical procedure estimation of technical efficiency in order to compare between the DMUs. Finally we conclude by a summary of the results with some remarks and open issues.

4.2 Efficiency measures for panel data

The principle of the efficiency measure estimation for panel data is the same as that for cross-sectional data. We need however to make additional assumptions about the temporal pattern of inefficiency. There are also differences between the two procedures in terms of the simulated likelihood function definition.

When information about all DMUs is available at T different time periods, it is preferable to use a stochastic frontier model which is adequate for panel data. In frontier analysis this model is more appropriate because even if it is not fundamentally different from the cross-sectional model, it has several advantages as it

- a. Increases the degree of freedom to estimate parameters;
- b. Provides consistent efficiency estimates when T is increasing;
- c. Does not require that the inefficiencies are independent of the regressors;

4.2.1 The panel data production frontier model

The panel stochastic frontier model, when the inefficiencies are assumed to vary systematically with time, is specified as follows

$$y_{it} = f(\underline{x}_{it}, \beta) + \epsilon_{it} = f(\underline{x}_{it}, \beta) + v_{it} - u_{it}, i = 1, \dots, n; t = 1, \dots, T \quad (4.2.1)$$

where $y_{it} = \log(Y_{it})$; Y_{it} : the observed output for observation i at the t^{th} time period (one output), so $Y_{it} \in \mathbb{R}_+$; $\underline{x}_{it} = \log(\underline{X}_{it})$; $\underline{X}_{it}$: a vector of length p describes the observed inputs for observation i at t , so $\underline{X}_{it} \in \mathbb{R}_+^p$ where p is the number of the inputs; β : a vector of unknown parameters to be estimated, $\beta \in \mathbb{R}^{l+(1 \times T)}$ where l is the number of parameters excluding the time-varying intercepts. If intercepts are constant over time, then $\beta \in \mathbb{R}^{l+1}$. Moreover, ϵ_{it} is the error term for observation i at time t , $f(\underline{x}_{it}, \beta)$ the production frontier, n is the number of DMUs under study and T is the number of periods or the number of observations for each DMU.

The two components of the error term are motivated by the idea that deviations from the frontier might not be entirely under the control of the DMU and that the performance of a DMU is affected by these two components. Hence, the term ϵ_{it} is divided into two parts, the inefficiency term u_{it} which is constrained to be non-negative ($u_{it} \geq 0$) and the statistical noise term v_{it} which is usually a normal with zero as mean and σ_V as standard deviation ($v_{it} \sim N(0, \sigma_V^2)$).

Furthermore, distributional assumptions will be imposed on both terms u_{it} and v_{it} . In particular, it is assumed that the components of the first are independently and identically positively distributed and the components of the second are independently and identically normally distributed. In addition, both terms are assumed continuous and independent of x_{it} . At first, it is supposed that the two terms are mutually independent and the model is estimated by the Maximum Likelihood (ML) method. At a second stage, they will be allowed to be dependent and the ML estimates of the first stage will be considered as initial values in the numerical optimization.

As proposed in the literature on panel SFA, the frontier model considers either the time-constant (see the Appendix C) or the time-varying efficiency. In our study we consider the frontier model with time-varying efficiency which is,

in our opinion, more realistic and reflects the inefficiency variability over time. Nevertheless, we shall limit our analysis to the fixed intercept over time and to a comparison between some time-varying models in order to select one of them.

4.2.2 Time-varying efficiency

When T is large, the assumption of a time-constant inefficiency is typically not appealing, as one would expect that inefficient DMUs are forced to improve over time. So, a time varying inefficiency is needed and a random-effects model should be used. To define a random-effects model, one has developed an extension of the fixed-effects model to a more general model to get consistent estimators of u_i when T is large. Among these we refer to Jondrow et al. (1982), which derived panel generalizations of the conditional inefficiency predictors, Battese and Coelli (1988) where the term u_i has a more general truncated-normal distribution, and Battese, Coelli and Colby (1989) which extend the model to allow unbalanced data.

The frontier model is called a random-effects model when it is described by

$$y_{it} = \beta_{0t} + \sum_{j=1}^l \beta_j x_{ijt} + v_{it} - u_{it} = \beta_{it} + \sum_{j=1}^l \beta_j x_{ijt} + v_{it} \quad (4.2.2)$$

where $\beta_{it} = \beta_{0t} - u_{it}$ is the intercept for DMU_i in the time period t and where β_{0t} is the intercept common to all $DMUs$ in the time period t , see e.g. Kumbhakar (1990) and Kumbhakar and Lovel (2000). Of course, $n \times T$ parameters β_{it} should be estimated but Cornwell, Schmidt and Sickles (1990) reduce this number to $3 \times n$ using

$$\beta_{it} = \Omega_{i1} + \Omega_{i2}t + \Omega_{i3}t^2 \quad (4.2.3)$$

where $\Omega_{ik}, k = 1, \dots, 3$ are parameters to be estimated. If $\Omega_{i2} = \Omega_{i3} = 0$, the model collapses to the time-constant efficiency model.

In the same way, Kumbhakar (1990) suggested a model in which the u_{it} are specified by following expression:

$$u_{it} = \eta(t) \cdot u_i = [1 + \exp \{ \eta_1 t + \eta_2 t^2 \}]^{-1} u_i \quad (4.2.4)$$

where u_i is assumed to have a half-normal distribution and η_1 and η_2 are two scalar parameters to be estimated. He suggested estimating the model with the maximum likelihood method but does not provide an empirical application. Battese and Coelli (1992) suggested a time-varying model for unbalanced panel data with an exponential function of time for u_{it} such as

$$u_{it} = \eta(t)u_i = [\exp \{-\eta_1 (t - T)\}] u_i \quad (4.2.5)$$

where u_i is assumed to have a truncated-normal distribution and η_1 is a scalar parameter to be estimated. They also proposed in their later work, Battese and Coelli (1995), a model where u_{it} follows a normal distribution truncated at zero. The ML estimation and the efficiency calculations of these cases have been included in the *FRONTIER* programs implemented by Coelli (1996).

Schmidt and Sickles (1984) suggested not to specify an implicit distribution for the inefficiency when the panel data are available and to estimate the fixed-effects model with the traditional panel data methods. In extension of this approach, Cornwell, Schmidt and Sickles (1990) and Lee and Schmidt (1993) have developed an approach in which they introduce the variation of the effect of inefficiencies over time. Both of the latter approaches propose a variation of inefficiencies more flexible than proposed in (4.2.4) and (4.2.5).

In this research, the Kumbhakar (1990) and the Cornwell, Schmidt and Sickles (1990) random-effects models will not be considered given the large number of parameters to be estimated, and so just one fixed intercept will be estimated. The Kumbhakar (1990), Battese and Coelli (1992) and a variety of time effects models are considered with a function $\eta(t) \geq 0$ that describes the evolution of inefficiency over time. These models are denoted as

$$M_1 : u_{it} = [\exp \{-\eta_1 (t - T)\}] u_i, \quad u_i \sim N^+ (0, \sigma_U^2); \quad (4.2.6)$$

$$M_2 : u_{it} = [\exp \{-\eta_1 (t - T)\}] u_i, \quad (4.2.7)$$

$$M_3 : u_{it} = [1 + \exp \{\eta_1 t + \eta_2 t^2\}]^{-1} u_i, \quad (4.2.8)$$

$$M_4 : u_{it} = [1 + \eta_1 (t - T) \sin (t - T)] u_i, ; \quad (4.2.9)$$

$$M_5 : u_{it} = [1 + \eta_1 \sin (\eta_2 t)] u_i, \quad (4.2.10)$$

$$M_6 : u_{it} = [1 + \eta_1 \sin (\eta_2 (t - T))] u_i, \quad (4.2.11)$$

$$M_7 : u_{it} = [1 + \eta_1 (t - T) \sin (\eta_2 (t - T))] u_i, \quad (4.2.12)$$

$$M_8 : u_{it} = \left[\eta_0 + \eta_1 t + \frac{1}{2} \eta_2 t^2 + 2 \sum_{h=1}^H (a_h \sin (ht) - b_h \cos (ht)) \right] u_i; \quad (4.2.13)$$

where, for all models and except for M_1 , $u_i \sim N^+(\mu, \sigma_U^2)$ and μ is the mean of the original normal distribution. That indicates that the inefficiency term u_i is a normal truncated at zero with mean μ . Furthermore, M_1 and M_2 are the Battese and Coelli (1992, 1995) models and M_3 is the Kumbhakar (1990) model. The proposed $M_4 - M_7$ models include sinusoidal functions to allow for possible periodicity effects in the inefficiency. For example, M_7 models a time-varying amplitude of the sine function depending on the parameter η_1 . The last considered model M_8 is the Fourier Flexible Form of Gallant (1984) which can closely approximate any smooth function $\eta(t)$ for sufficiently large H . In our study the Akaike Information Criterion (AIC) is used to select a model among M1 to M8. Moreover, a likelihood ratio test is performed for the nested models.

4.2.3 Model estimation

Considering the models described by (4.2.2) and (4.2.6)-(4.2.13), the methods used to estimate all models depend on the distributional assumptions. When v_{it} is i.i.d. normal, u_{it} is i.i.d. with positive support, v_{it} and u_{it} are mutually independent and independent from the regressors, the Maximum Likelihood Estimation (MLE) is feasible. Schmidt and Sickles (1984) conjectures that given suitable regularity conditions the ML estimates of (4.2.2) with (4.2.6)

are consistent and asymptotically efficient as $n \rightarrow \infty$ regardless of T . In the particular case where $v_{it} \sim iidN(0, \sigma_V^2)$ and $u_{it} \sim iidN^+(0, \sigma_U^2)$, the MLE leads for $\epsilon_i = (\epsilon_{i1}, \dots, \epsilon_{iT})'$ to the log-likelihood function

$$\begin{aligned} l = \ln(L) &= cte - \frac{n}{2} \ln \sigma_*^2 - \frac{1}{2} \sum_i^n a_{*i} - \frac{n.T}{2} \ln \sigma_V^2 - \frac{n}{2} \ln \sigma_U^2 \\ &\quad + \sum_i^n \ln \left[1 - \Phi \left(-\frac{\mu_{*i}}{\sigma_*} \right) \right], \end{aligned} \quad (4.2.14)$$

which leads to the technical efficiency (TE) estimate for all $i = 1, \dots, n$

$$\begin{aligned} TE_{it} &= E(\exp\{-u_{it}\} \mid \epsilon_i) \\ &= \frac{1 - \Phi\left(\eta(t) \sigma_* - \frac{\mu_{*i}}{\sigma_*}\right)}{1 - \Phi\left(-\frac{\mu_{*i}}{\sigma_*}\right)} \exp\left\{-\eta(t) \mu_{*i} + \frac{1}{2} \eta^2(t) \sigma_*^2\right\}, \end{aligned} \quad (4.2.15)$$

where L is the likelihood function, cte is a scalar, $\eta(t)$ is the function in (4.2.5), $\sigma_*^2 = \frac{\sigma_V^2 \sigma_U^2}{\sigma_V^2 + \sigma_U^2 \sum_t \eta^2(t)}$, $\mu_{*i} = \frac{(\sum_t \eta(t) \epsilon_{it}) \sigma_V^2}{\sigma_V^2 + \sigma_U^2 \sum_t \eta^2(t)}$, $a_{*i} = \frac{1}{\sigma_V^2} \left[\sum_t \epsilon_{it}^2 - \frac{\sigma_U^2 (\sum_t \eta(t) \epsilon_{it})^2}{\sigma_V^2 + \sigma_U^2 \sum_t \eta^2(t)} \right]$ and Φ is the standard normal cdf. When $\eta(t)$ is replaced by the expression with two unknown parameters given in (4.2.4), we obtain the TE estimate of the model described by (4.2.2) and (4.2.4). See the Appendix C for more details.

In comparison with the cross section data, calculation of the log-likelihood function for panel data is similar, but it changes at the level of computing the ϵ_i density which is $g(\epsilon_i)$ and where $\epsilon_i = (\epsilon_{i1}, \dots, \epsilon_{it}, \dots, \epsilon_{iT})'$. In the expression of $g(\epsilon_i)$, the joint density $f(u_i, v_i)$ of u_i and v_i is replaced by $f_1(u_i) f_2(v_i) = f_1(u_i) \prod_t f_2(\epsilon_{it} + \eta(t) u_i)$. Of course, this last expression is integrated by u_i to get $g(\epsilon_i)$.

When u_i and v_i are dependent, the joint density of them when panel data is available becomes

$$\begin{aligned} f_1(u_i) f_2(v_i) c_\theta(F_1(u_i), F_2(v_i)) &= \\ f_1(u_i) \prod_t f_2(\epsilon_{it} + \eta(t) u_i) \prod_t c_\theta(F_1(u_i), F_2(\epsilon_{it} + \eta(t) u_i)), \end{aligned} \quad (4.2.16)$$

where c is a bivariate copula density which expresses the dependence between the two variables u_i and v_i , and $F_1(u_i)$ and $F_2(v_i)$ are two uniform variables which are the cdf of $f_1(u_i)$ and $f_2(v_i)$ respectively and called the margins. The independence case is a special case of this model when the copula is the product copula, for which $c(\cdot, \cdot) = 1$. But for general copula functions, the ML estimation will become more complicated.

Given that the v_{it} are supposed independent and identically distributed, the density of ϵ_i becomes

$$g(\epsilon_i) = \int_0^\infty f(\epsilon_i, u_i) du_i \quad (4.2.17)$$

$$\begin{aligned} &= \int_0^\infty f(\epsilon_{i1}, \dots, \epsilon_{it}, \dots, \epsilon_{iT}, u_i) du_i \\ &= \int_0^\infty f_1(u_i) \prod_t f_2(\epsilon_{it} + \eta(t) u_i) \\ &\quad c_\theta(F_1(u_i), F_2(\epsilon_{i1} + \eta(1) u_i), \dots, F_2(\epsilon_{iT} + \eta(T) u_i)) du_i \\ &= \int_0^\infty f_1(u_i) \prod_t f_2(\epsilon_{it} + \eta(t) u_i) \\ &\quad \prod_t c_\theta(F_1(u_i), F_2(\epsilon_{it} + \eta(t) u_i)) du_i \\ &= \int_0^\infty f_1(u_i) \\ &\quad \prod_t \left[f_2(\epsilon_{it} + \eta(t) u_i) c_\theta(F_1(u_i), F_2(\epsilon_{it} + \eta(t) u_i)) \right] du_i, \\ &= \int_0^\infty f_1(u_i) \prod_t A_{it} du_i = E \left(\prod_t A_{it} \right), \end{aligned} \quad (4.2.18)$$

where $A_{it} = f_2(\epsilon_{it} + \eta(t) u_i) c_\theta(F_1(u_i), F_2(\epsilon_{it} + \eta(t) u_i))$. Therefore, assuming the independence across DMUs, the log-likelihood function can be written as

$$\begin{aligned} l(\vartheta) &= \log L(\vartheta) \\ &= \log L(\sigma_U, \sigma_V, \theta, \beta_0, \beta, \underline{\eta}_k) \\ &= \sum_{i=1}^n \log g_\theta(\epsilon_i) = \sum_{i=1}^n \log g_\theta \left(y_i - \left(\beta_0 + \sum_{j=1}^l \beta_j x_{ij} \right) \right), \end{aligned} \quad (4.2.19)$$

where $\beta_0 = (\beta_{01}, \dots, \beta_{0t}, \dots, \beta_{0T})'$ and $\beta = (\beta_1, \dots, \beta_j, \dots, \beta_l)'$ are vectors with a length equal respectively to the time periods T and the number of inputs l , $\underline{\eta}_k$ is a vector of k parameters in the time-varying function and where $y_i = (y_{i1}, \dots, y_{it}, \dots, y_{iT})'$ and $x_{ij} = (x_{ij1}, \dots, x_{ijt}, \dots, x_{ijT})$. For simplicity, all intercepts β_{0t} , $t = 1, \dots, T$ are considered the same and denoted by β_0 in the empirical analysis.

Generally, the expression of the function $l(\vartheta)$ is complex in the dependence case and to obtain analytical derivatives becomes a tedious or even impossible task in several cases. So, the log-likelihood is optimized numerically using the *mle* function in the R software and using the simplex numerical method called the Nelder-Mead method.

Once the parameters $(\sigma_U, \sigma_V, \theta, \beta_0, \beta, \underline{\eta}_k)$ are estimated, the technical efficiency can be estimated using the expected value of $(\exp \{-u_{it}\} \mid \epsilon_i)$ as

$$\begin{aligned}
 TE_{it} &= E\left(\exp \{-u_{it}\} \mid \epsilon_i\right) & (4.2.20) \\
 &= E\left(\exp \{-\eta(t) u_i\} \mid (\epsilon_{i1}, \dots, \epsilon_{it}, \dots, \epsilon_{iT})\right) \\
 &= \int_0^{+\infty} \exp \{-\eta(t) u_i\} f_1(u_i \mid (\epsilon_{i1}, \dots, \epsilon_{it}, \dots, \epsilon_{iT})) du_i \\
 &= \int_0^{+\infty} \exp \{-\eta(t) u_i\} \frac{f(u_i, \epsilon_i)}{g(\epsilon_i)} du_i \\
 &= \frac{1}{g(\epsilon_i)} \int_0^{+\infty} f_1(u_i) \exp \{-\eta(t) u_i\} \prod_t A_{it} du_i \\
 &= \frac{E\left(\exp \{-\eta(t) u_i\} \prod_t A_{it}\right)}{E\left(\prod_t A_{it}\right)}. & (4.2.21)
 \end{aligned}$$

Given again the complexity of the TE_{it} expression, the expectation will be estimated for a large number m of Monte Carlo draws by

$$\widehat{TE_{it}} \cong \frac{\frac{1}{m} \sum_{j=1}^m \left(\exp \{-\eta(t) u_j\} \prod_t A_{ijt} \right)}{\frac{1}{m} \sum_{j=1}^m \left(\prod_t A_{ijt} \right)}. \quad (4.2.22)$$

4.3 Inferences on the Technical Efficiency measure

Since TE of each DMU at each time t is unknown and it is estimated by $\widehat{TE}$, an inference about it is required. To build the confidence interval at a level α , given that the true sampling distribution is not available, we see that a modified algorithm of the parametric bootstrap Algorithm#3 of Simar and Wilson (2010) adapted to the dependence case and to the panel framework is more appropriate. Hence, we developed a procedure to estimate the associated confidence bounds when v is normal and u is half-normal which can be generalized for any positive function of u such as the truncated-normal one. The method is easy to apply but it is quite computationally intensive. The steps are the following:

1. Estimate $\vartheta = (\sigma_U, \sigma_V, \theta, \beta_0, \beta, \underline{\eta}_k)$ according to (4.2.19), using the observed $(x_{it}, y_{it}), i = 1, \dots, n$ and $t = 1, \dots, T$ and using a numerical optimization procedure to get $\widehat{\vartheta} = (\widehat{\sigma}_U, \widehat{\sigma}_V, \widehat{\theta}, \widehat{\beta}_0, \widehat{\beta}, \widehat{\underline{\eta}}_k)$ and to compute the point estimates $\widehat{TE}$ as described before.
2. For $i = 1, \dots, n$, draw $u_i^* \sim N^+(0, \widehat{\sigma}_U^2)$ and $v_{it}^* \sim N(0, \widehat{\sigma}_V^2), t = 1, \dots, T$ such that u_i^* and v_{it}^* are dependent with dependence characterized by the Clayton copula. Then compute $y_{it}^* = \widehat{\beta}_0 + \sum_{j=1}^l \widehat{\beta}_j x_{ijt} + v_{it}^* - \widehat{\eta}(t) u_i^*$. There are several procedures to generate the pair (u_i^*, v_{it}^*) according to the Clayton copula, we use the one described in Nelsen (1999), page 41. The four steps of this procedure are
 - a. Draw $T+1$ independent uniform random variables $w_{1i}, h_{2i1}, \dots, h_{2it}, \dots, h_{2iT}$, such that $w_{1i} \sim U(0, 1)$ and $h_{2it} \sim U(0, 1)$ for $t = 1, \dots, T$.
 - b. Set $w_{2it} = \left[w_{1i}^{-\widehat{\theta}} \left(h_{2it}^{-\widehat{\theta}/(1+\widehat{\theta})} - 1 \right) + 1 \right]^{-1/\widehat{\theta}}$, for all $t = 1, \dots, T$.
 - c. Set $u_i^* = F_1^{-1}(w_{1i})$ and $v_{it}^* = F_2^{-1}(w_{2it})$, for all $t = 1, \dots, T$ and where F_1 and F_2 are the cdf of the $N^+(0, \widehat{\sigma}_U^2)$ and $N(0, \widehat{\sigma}_V^2)$ respectively.
 - d. Repeat steps a to c n times to generate $n \times T$ pairs (u_i^*, v_{it}^*) .

3. Using the pseudo-data $\mathcal{S}_{b,n}^* = \{(\underline{x}_{it}, y_{it}^*)\}_{i=1}^n$, compute bootstrap estimates $\hat{\vartheta}_b^* = \arg \max_{\vartheta \in \Theta} l(\vartheta \mid \mathcal{S}_{b,n}^*)$ after replacing y_{it} by y_{it}^* in (4.2.19) and then compute the bootstrap estimates $\widehat{TE}_b^*$ using (4.2.21) after replacing ϵ by $\epsilon_b^* = y - \hat{\beta}_0^* - \hat{\beta}^* \cdot x$, where $\underline{x}_{it}$ and y_{it} represent the observed data.

4. Repeat steps 2 and 3, B times to obtain estimates $\mathcal{B}^* = \{\hat{\vartheta}_b^*\}_{b=1}^B$. Therefore, use $\mathcal{B}^*$ to get $\mathcal{E}^* = \{\widehat{TE}_b^*\}_{b=1}^B$. Each individual i is described by a sub-matrix of $\mathcal{E}^*$ denoted $\mathcal{E}_i^*$, it has T rows and B columns.

For each individual i at time period t (row t of the $\mathcal{E}_i^*$ matrix, denoted $\mathcal{E}_{it}^*$), $i = 1, \dots, n$, compute the $(\frac{\alpha}{2})$ and the $(1 - \frac{\alpha}{2})$ quantiles for $\mathcal{E}_{it}^*$ by considering its B components. The $100 \times (1 - \alpha)$ percentile bootstrap confidence interval of the statistic of interest TE is obtained by the probability $P\left((\mathcal{E}_{it}^*)_{\frac{\alpha}{2}} < TE_{it} < (\mathcal{E}_{it}^*)_{1-\frac{\alpha}{2}}\right) = 1 - \alpha$.

Hence, using the $100 \times (\frac{\alpha}{2})$ and $100 \times (1 - \frac{\alpha}{2})$ percentiles, we define the lower and the upper bounds of the confidence interval as $TE_{it} \in \left[(\mathcal{E}_{it}^*)_{\frac{\alpha}{2}}, (\mathcal{E}_{it}^*)_{1-\frac{\alpha}{2}}\right]$.

We note that the estimation procedure presented in Section 4.2.3 leads sometimes to a positive skewness of the composite error term which consequently leads to biased parameter estimates and to biased technical efficiencies estimates because all of these latter will be close to one. If this is the case, the procedure presented in this section allows us to overcome this problem.

To perform our procedure, a simulation example is proposed. The model describing data is supposed to be log-linear where there are one input and one output such that for all $i = 1, 2, \dots, n$ and for all $t = 1, 2, \dots, T$, we have $\log(Y_{it}) = \beta_0 + \beta_1 \log(10(1 + X_{it}))$ where $X_{it} \sim U(0, 1)$ and parameters will be set to $\beta_0 = \log(10)$ and $\beta_1 = 1$. As for the noise term and the inefficiency term, they are supposed to be normal as usual for the first such that $v_{it} \sim N(0, \sigma_V^2)$ with $\sigma_V = 0.5$ and half-normal for the second such that $u_{it} = \eta(t)u_i$ and $u_i \sim N^+(0, \sigma_U^2)$ with $\sigma_U = 1$ and the two components are dependent using the Clayton copula with dependence parameter $\theta = 1$. The time varying function is supposed to be $\eta(t) = \exp\{-\eta_1 \cdot (t - T)\}$ with $\eta_1 = -0.1$. We suppose that $n = 50$, $T = 10$ and $B = 500$. To compute the true efficiencies, the number of simulations to approximate numerically the integral is set to $m = 10000$ which

is large enough to have a good approximation of the expectation in equation (4.2.21) evaluated at the true values.

The bootstrap procedure shows that all estimated efficiencies are covered by their confidence intervals. The percentage for the true efficiencies is evaluated at 87%, 91.2% and 100% for a significance level of 10%, 5% and 1% respectively, so that the bootstrap coverage ratio is reasonably close to the nominal level given our moderate number of bootstrap replications.

4.4 Illustrative empirical analysis on Moroccan drinking water area

The data set analyzed in this section and chosen to illustrate our methodology contains information on the water production and its sales to the subscribers and to the municipal utilities called the self-governance in Morocco. The national public company called the National Office of the Drinking Water (ONEP according to its French abbreviation) ensures the largest part of the production, the pipe and the distribution of water in the entire national territory. It produces more than 80 percent of the country's drinking water. This sector depends mainly on the domestic consumption. So, to strengthen water resources and to rationalize its use, starting in the 1980s it led to certain administrative and technical actions such as an information campaign and the installation of individual water meters in households, in order for example to reduce wasting.

We are interested in this practical case in the efficiency of certain national participants in the management of this particular and vital good. To do so, the Farrel technical efficiency rate will be estimated using a panel data set in order to compare the performance of certain Moroccan provinces with respect to their produced quantities in the sector and to their water sales. We shall analyze, thus, the degree of efficiency of every producing entity of water in order to situate it among the others at the national level and we shall know how much should be its sales to attain efficiency. Efficiency being an estimate, we also provide confidence intervals.

The considered variables in this study are the ONEP's sales and the number

of the subscribers as inputs and the water production as output. Both of sales and production are evaluated in thousand cubic meter ($1000 m^3$). Therefore, the model will be a simple model with two independent variables and the frontier function chosen to describe the production technology is the translog function. As for the copula function, the Archimedean Clayton copula is used because it is popular in empirical applications, it is flexible and easy to construct, and it nests the independence copula as a special case, see for example Bhat and Eluru (2009).

Hence, the data represent sales, the number of subscribers and the water production for a set of 50 provinces through 15 Moroccan regions for a duration of seven years from 2001 until 2007. Six provinces among a total of 56 were omitted because of lack or unavailability of data as indicated in the statistical yearbooks of the corresponding years published annually by the Statistics Direction and which have as source the ONEP entity in this area. These six provinces are Tan Tan, Inezgane-Ait-Melloul, Sidi Youssef Ben Ali and Al Ismailia in respectively Guelmim-Es-Semara, Souss-Massa-Daraâ, Marrakech-Tensift-Al Haouz and Meknès-Tafilalet regions and all provinces of Grand-Casablanca region which are Casablanca and Mohammedia.

Being complex, the optimization of the log-likelihood function is performed using numerical optimization in three steps:

- Step 1. The model (4.2.1) is fitted using the pooled-OLS regression² without the technical inefficiency term. Hence, the inefficiency is null ($u_{it} = 0$) and the time effect is zero.
- Step 2. The OLS parameters of step 1, except the intercept β_0 which is biased, are used as initial values in this step to estimate numerically the model (4.2.1) using the maximum likelihood estimation (MLE) assuming independence between u and v . The β_0 parameter is adjusted by shifting it according to the Corrected Ordinary Least Squares procedure (COLS) used in the R `frontier` package, see e.g. Coelli (1995a) .

- Step 3. In this last step, MLE numerical optimization is performed using the

²The pooled-OLS: Treat all the observations for all time periods in the panel data as a single sample and use the OLS to fit the model.

step 2 estimates as initial values. As for the initial value of the copula parameter θ , a grid of values for θ is given. Given that the models are the same (the difference is just the value of θ), the comparison of the log-likelihood function is done directly (without estimation) and hence the value which gives the highest log-likelihood value is chosen as initial value of θ in this step.

About the model, the full translog function with two exogenous variables is considered in this analysis. Inputs are the number of subscribers and the total of sales. So, we will have the following number of parameters: one for the intercept, two for the variables, two for their terms squared and one for their interaction. The others are σ_U , σ_V , θ and $\underline{\eta}_k$ and μ is added when the truncated-normal distribution is considered. The estimation of the model with the full translog function in the case of the independence, using the **frontier** package of the R software, has revealed that the number of subscribers variable, its squared term and the interaction term are not significant and consequently we do not reject that their coefficients are equal to zero. For this reason, only the sales and its square term will be included in the model. Hence, the total number of parameters included in the final considered model in the dependence case is at least seven parameters.

Furthermore, the full translog model is not used because the minimal *AIC* criterion which is $-2\log Lik + 2k$, where k is the number of parameters to be estimated in the model, is smaller for the restricted function evaluated at -78.0919 in comparison with the full function evaluated at -75.3816; and being nested, the likelihood ratio test rejects the full model in favor of the restricted one at the 5% level of significance. Estimates are used as initial values in the numerical optimization in the case where v_{it} is normal and u_{it} is half-normal or truncated-normal when dependence between them is considered.

According to the expressions of the time-varying function and to the inefficiency distribution, Panel Stochastic Frontier models described by (4.2.1) and (4.2.6)-(4.2.13) are estimated and the model M_7 which has the time-varying expression (4.2.12) is selected according the minimum *AIC* criterion with a log-likelihood value evaluated at 595.041 as pointed out in Table 4.1. In addition, M_4 and M_7 being nested, the latter one was not rejected according to the

Table 4.1: The AIC values of the estimated models

Model	M_1	M_2	M_3	M_4	M_5	M_6	M_7	M_8
$\log L$	413.384	418.518	400.134	502.000	433.728	492.155	595.041	165.156
df	7	8	9	8	9	9	9	14
AIC	-812.768	-821.036	-782.268	-988.000	-849.456	-966.310	-1172.082	-302.312

Table 4.2: The model estimation with correlated error terms

	Estimate	Std. Error	t value	$Pr(T > t)$
σ_U	3.2313	0.000400	8084.890	< 1e-16
σ_V	0.0539	0.000215	250.615	< 1e-16
β_0	3.7094	0.001028	3607.835	< 1e-16
β_1	1.6721	0.000161	10413.472	< 1e-16
β_2	-0.0858	0.000082	-1048.721	< 1e-16
η_1	0.0425	0.000668	63.676	< 1e-16
η_2	0.2276	0.000625	364.036	< 1e-16
θ	2.0389	0.000352	5789.226	< 1e-16
μ	0.0179	0.000607	29.452	< 1e-16
$-2\log L$	-1190.083			

likelihood ratio test.

First of all, it is noted that for this chosen model the time effect is significant, μ is not equal to zero and the two terms of inefficiency u and noise v are dependent. The parameter θ is not close to zero and hence the Clayton copula does not approach the Product copula related to the independence of the two error terms. All other parameters are statistically significant in the model as pointed out in Table 4.2.

As for the efficiency scores, which are one of our major objectives, Table 4.3 presents the estimation results of the model (4.2.1) under assumptions (4.2.12) denoted M_7 and where v_{it} and u_{it} are correlated by the Clayton copula. It mainly shows that all provinces are technically inefficient. The most efficient are Rabat-Skhirate-Témara and El Jadida in respectively Rabat-Salé-Zemmour-Zaer and Doukkala-Abda regions with a mean score for the period greater than 0.65 and the most inefficient one is Ben Slimane province in Chaouia-Ouadigha region. Rabat-Skhirate-Témara is near the frontier and should on average increase its sales by just about 0.3% to be efficient. Moreover, Table 4.4 which summarizes the previous one shows also that only 20% of the provinces exceed the national efficiency mean for the entire period evaluated at 0.102 which is a very weak score. Even if the average is low, it has progressed but slowly from

year to the next one which may indicate that the ONEP policy in the production and selling of water was not adequate during the seven studied years. Moreover, the efficiency standard deviation is large because it is estimated at 0.176. Hence, with the exception of one or two provinces, scores are very low and the dispersion is high which reflect the mediocre performance of the sector with respect to the relation between the ONEP's drinking water production and its sales.

It is clear also that the time effect on the efficiency scores is positive given that $\hat{\eta}_1$ and $\hat{\eta}_2$ are positive and the TE increases over the period under study as specified previously. Even if this effect is weak, it is statistically highly significant with a p-value less than 1e-16.

Table 4.3: Technical Efficiency scores for the 2001-2007 period and their confidence intervals

Nr	DMUs Name	$\widehat{TE}_{i1}$	$\widehat{TE}_{i2}$	$\widehat{TE}_{i3}$	$\widehat{TE}_{i4}$	$\widehat{TE}_{i5}$	$\widehat{TE}_{i6}$	$\widehat{TE}_{i7}$	$\widehat{TE}_i$	$\widehat{bias}_i^*$	$\widehat{\sigma}_i^*$	$\widehat{TE}_{cor_i}$	Lower	Upper
1	Oued Ed-Dahab	0.0143	0.0173	0.0211	0.0254	0.0294	0.0323	0.0333	0.0247	-0.0019	0.0118	0.0264	0.0084	0.0506
2	Boujdour	0.0057	0.0072	0.0092	0.0115	0.0137	0.0154	0.0160	0.0112	-0.0003	0.0039	0.0115	0.0051	0.0195
3	Laâyoune	0.0293	0.0344	0.0406	0.0473	0.0534	0.0578	0.0594	0.0460	-0.0054	0.0275	0.0517	0.0128	0.1105
4	Assa-Zag	0.0057	0.0072	0.0092	0.0114	0.0137	0.0153	0.0160	0.0112	0.0006	0.0040	0.0108	0.0057	0.0203
5	Es-Semara	0.0086	0.0107	0.0134	0.0164	0.0193	0.0215	0.0223	0.0160	0.0001	0.0071	0.0160	0.0067	0.0326
6	Guelmim	0.0297	0.0348	0.0411	0.0478	0.0540	0.0583	0.0599	0.0465	-0.0045	0.0253	0.0508	0.0143	0.1038
7	Tata	0.0096	0.0119	0.0148	0.0181	0.0212	0.0235	0.0244	0.0176	-0.0005	0.0073	0.0183	0.0072	0.0337
8	Agadir-Ida ou Tanane	0.1623	0.1762	0.1920	0.2076	0.2210	0.2302	0.2334	0.2032	-0.0280	0.2269	0.2032	0.0364	0.8463
9	Chtouka-Ait Baha	0.0097	0.0119	0.0148	0.0181	0.0212	0.0236	0.0244	0.0177	0.0001	0.0080	0.0176	0.0072	0.0361
10	Ouarzazate	0.0301	0.0352	0.0415	0.0483	0.0545	0.0589	0.0606	0.0470	-0.0021	0.0314	0.0488	0.0142	0.1251
11	Taroudannt	0.0271	0.0319	0.0378	0.0441	0.0500	0.0541	0.0557	0.0429	-0.0028	0.0269	0.0461	0.0127	0.1086
12	Tiznit	0.0179	0.0215	0.0259	0.0308	0.0354	0.0388	0.0400	0.0300	0.0001	0.0186	0.0299	0.0101	0.0762
13	Zagora	0.0145	0.0176	0.0214	0.0257	0.0298	0.0327	0.0338	0.0251	-0.0010	0.0120	0.0262	0.0092	0.0518
14	Kenitra	0.1200	0.1322	0.1460	0.1599	0.1721	0.1804	0.1833	0.1563	-0.0160	0.1420	0.1614	0.0355	0.5265
15	Sidi Kacem	0.0671	0.0759	0.0861	0.0967	0.1062	0.1128	0.1151	0.0943	-0.0087	0.0657	0.1007	0.0246	0.2559
16	Ben Slimane	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000	0.0008	0.0000	0.0000	0.0025
17	Khouribga	0.0199	0.0238	0.0286	0.0339	0.0387	0.0423	0.0436	0.0330	-0.0045	0.0238	0.0375	0.0078	0.0920
18	Settat	0.0057	0.0072	0.0092	0.0115	0.0137	0.0154	0.0160	0.0113	-0.0012	0.0094	0.0125	0.0025	0.0355
19	Al Haouz	0.0111	0.0137	0.0169	0.0205	0.0239	0.0264	0.0274	0.0200	-0.0002	0.0091	0.0205	0.0078	0.0406
20	Chichaoua	0.0095	0.0118	0.0146	0.0179	0.0210	0.0233	0.0241	0.0175	-0.0002	0.0075	0.0176	0.0072	0.0343
21	El Kelaâ des Sraghna	0.0314	0.0368	0.0433	0.0502	0.0566	0.0611	0.0627	0.0489	-0.0034	0.0317	0.0525	0.0140	0.1264
22	Essaouira	0.0204	0.0243	0.0292	0.0346	0.0395	0.0431	0.0444	0.0336	-0.0019	0.0193	0.0361	0.0106	0.0791
23	Marrakech Ménara	0.3315	0.3485	0.3671	0.3849	0.3999	0.4098	0.4133	0.3793	-0.0585	0.3183	0.4022	0.0582	0.9839
24	Berkane - Taourirt	0.0820	0.0919	0.1033	0.1151	0.1255	0.1326	0.1352	0.1122	-0.0110	0.0949	0.1198	0.0262	0.3574
25	Figui	0.0108	0.0133	0.0164	0.0199	0.0233	0.0258	0.0267	0.0195	-0.0021	0.0073	0.0216	0.0073	0.0336
26	Jerada	0.0135	0.0165	0.0202	0.0243	0.0281	0.0310	0.0320	0.0236	0.0003	0.0119	0.0234	0.0094	0.0516
27	Nador	0.0655	0.0741	0.0842	0.0947	0.1041	0.1106	0.1129	0.0923	-0.0095	0.0750	0.0989	0.0218	0.2844
28	Oujda	0.0551	0.0628	0.0720	0.0816	0.0901	0.0961	0.0983	0.0794	-0.0063	0.0708	0.0848	0.0184	0.2692
29	Khemisset	0.0355	0.0413	0.0483	0.0558	0.0626	0.0674	0.0692	0.0543	-0.0026	0.0409	0.0577	0.0148	0.1594

Table 4.3: Technical Efficiency scores for the 2001-2007 period and their confidence intervals

Nr	DMUs Name	$\widehat{TE}_{i1}$	$\widehat{TE}_{i2}$	$\widehat{TE}_{i3}$	$\widehat{TE}_{i4}$	$\widehat{TE}_{i5}$	$\widehat{TE}_{i6}$	$\widehat{TE}_{i7}$	$\widehat{TE}_i$	$\widehat{bias}_i^*$	$\widehat{\sigma}_i^*$	$\widehat{TE}_{cor_i}$	Lower	Upper
30	Rabat-Skhirate-Ténara	0.9966	0.9968	0.9970	0.9971	0.9972	0.9973	0.9973	0.9970	-0.0419	0.3040	0.9970	0.1733	1.0000
31	El Jadida	0.6315	0.6448	0.6588	0.6720	0.6827	0.6898	0.6922	0.6674	-0.0966	0.3282	0.7533	0.1110	1.0000
32	Safi	0.0636	0.0721	0.0820	0.0924	0.1016	0.1080	0.1103	0.0900	-0.0077	0.0847	0.0931	0.0205	0.3187
33	Azilal	0.0215	0.0256	0.0307	0.0362	0.0413	0.0450	0.0463	0.0352	-0.0039	0.0170	0.0389	0.0113	0.0715
34	Beni Mellal	0.1289	0.1415	0.1558	0.1702	0.1826	0.1911	0.1942	0.1663	-0.0151	0.1671	0.1663	0.0362	0.6114
35	El Hajeb	0.0082	0.0103	0.0128	0.0158	0.0186	0.0207	0.0215	0.0154	-0.0008	0.0081	0.0163	0.0053	0.0341
36	Errachidia	0.0469	0.0539	0.0623	0.0710	0.0790	0.0845	0.0865	0.0692	-0.0069	0.0487	0.0762	0.0180	0.1906
37	Ifrane	0.0239	0.0283	0.0337	0.0396	0.0450	0.0489	0.0503	0.0385	-0.0026	0.0223	0.0415	0.0121	0.0909
38	Khénifra	0.0310	0.0362	0.0427	0.0496	0.0559	0.0604	0.0620	0.0482	-0.0045	0.0367	0.0535	0.0100	0.1385
39	Meknès El Menzeh	0.0452	0.0521	0.0602	0.0688	0.0765	0.0820	0.0840	0.0670	-0.0115	0.0495	0.0779	0.0000	0.1688
40	Boulmane	0.0109	0.0134	0.0165	0.0201	0.0234	0.0259	0.0268	0.0196	0.0001	0.0093	0.0191	0.0078	0.0412
41	Fès	0.3423	0.3594	0.3780	0.3958	0.4108	0.4207	0.4241	0.3902	-0.0587	0.3274	0.4020	0.0594	0.9844
42	Sefrou	0.0234	0.0277	0.0331	0.0389	0.0442	0.0481	0.0495	0.0378	-0.0014	0.0259	0.0392	0.0113	0.1032
43	Zouagha My Yacoub	0.0084	0.0104	0.0130	0.0160	0.0189	0.0210	0.0218	0.0156	-0.0005	0.0059	0.0161	0.0066	0.0281
44	Al Houcēma	0.0297	0.0349	0.0411	0.0479	0.0540	0.0584	0.0600	0.0466	-0.0056	0.0286	0.0523	0.0126	0.1144
45	Taounate	0.0304	0.0357	0.0420	0.0488	0.0551	0.0595	0.0611	0.0475	-0.0031	0.0252	0.0511	0.0156	0.1049
46	Taza	0.0279	0.0328	0.0388	0.0453	0.0512	0.0555	0.0570	0.0441	-0.0058	0.0245	0.0499	0.0123	0.0995
47	Chefchaouen	0.0199	0.0238	0.0286	0.0339	0.0388	0.0423	0.0436	0.0330	-0.0021	0.0169	0.0349	0.0114	0.0713
48	Larache	0.0567	0.0646	0.0739	0.0836	0.0923	0.0984	0.1006	0.0815	-0.0112	0.0657	0.0909	0.0179	0.2481
49	Tanger	0.2548	0.2711	0.2891	0.3066	0.3215	0.3314	0.3349	0.3013	-0.0365	0.2979	0.3013	0.0501	0.9696
50	Tétouan	0.1522	0.1658	0.1811	0.1964	0.2096	0.2185	0.2217	0.1922	-0.0275	0.2205	0.1922	0.0354	0.8186

where $\widehat{TE}_i$, $\widehat{bias}_i^*$, $\widehat{\sigma}_i^*$, $\widehat{TE}_{cor_i}$, $\widehat{Lower}$ and $\widehat{Upper}$ are the means according T of their corresponding expressions

TE_cor: Bias corrected efficiency

Table 4.4: Provinces number in each region according to the mean of TE estimates

Region Name	[0, 0.2[	[0.2, 0.4[	[0.4, 0.6[	[0.6, 0.8[	[0.8, 1[
Oued Ed-Dahab - Lagouira	1				
Laâyoune-Boujdour-S. EL Hamra	2				
Guelmim - Es-Semara	4				
Souss - Massa - daraâ	5	1			
Gharb - Chhrarda - Béni Hssen	2				
Chaouia - Ouardigha	3				
Marrakech - Tensift - Al Haouz	4	1			
Oriental	5				
Rabat-Salé-Zemmour-Zaer	1				1
Doukala-Abda	1			1	
Tadla - Azilal	2				
Meknès - Tafilalet	5				
Fès - Boulemane	3	1			
Taza - Al Hoceïma - Taounate	3				
Tanger - Tétouan	3	1			

The table is the same according to the median of TE estimates

Inference on the technical efficiency measure is made using a parametric percentile bootstrap procedure to estimate robust confidence intervals of the statistic of interest. Bootstrap samples are obtained according to the step 2 of the procedure in Section 4.3 and Table 4.3 presents technical efficiency estimates for each province and for the 2001-2007 period, their means, their corrected bias and the mean of the estimated lower and upper confidence interval bounds following the rest of the steps of the same bootstrap procedure. So, for each province, the mean of $\widehat{TE}_i$ is bounded by the mean of $(\mathcal{E}_{it}^*)_{\frac{\alpha}{2}}$ and the mean of $(\mathcal{E}_{it}^*)_{1-\frac{\alpha}{2}}$ for the seven years. Inference performed with $B = 500$ bootstrap replications shows that TE estimates for all provinces are in their corresponding confidence intervals at a 5% significant level but with a relatively large range for that which have great TE scores as depicted in Table 4.3. However, the global average width of the intervals is 0.22 with fifteen provinces (30% of all) having a width bigger than this average.

Technical efficiencies being estimates and in order to know if the bias correction is needed, their bias were estimated using the bootstrapped efficiencies as defined in Daraio and Simar (2007, p. 55) but using the median instead of the mean given that the TE distribution is skewed (mean estimated at 0.1024 is greater than median estimated at 0.0423). Then, for each individual i at time period t the bias is expressed by $\widehat{bias}^*(\widehat{TE}_{it}) = \widehat{median}(\widehat{TE}_b^*)_{it} - \widehat{TE}_{it}$ and

the bias corrected estimate of TE is defined as $\widehat{TE}_{cor_{it}} = \widehat{TE}_{it} - \widehat{bias}^*(\widehat{TE}_{it})$. Generally, the correction is needed if $|\widehat{bias}^*(\widehat{TE}_{it})|/\widehat{\sigma}_{it}^* > 0.25$ as pointed out in Efron(1982), where $\widehat{bias}^*(\widehat{TE}_{it})$ is the estimated bias of the bootstrap estimates and $\widehat{\sigma}_{it}^*$ is their standard error. Indeed, the bias was important for forty-four DMUs among fifty and the ratio reached on average its maximum with 1.03 points for the Assa-Zag province. The correction has reduced the average of the scores of the period for six DMUs and has increased this average for the thirty eight others. So, the efficiency scores were overestimated for the first ones and underestimated for the last ones. Furthermore, the means of the corrected TE are well in their confidence intervals.

4.5 Conclusion

This chapter presented and proposed, at first, the panel stochastic frontier analysis when the error terms are dependent and secondly an associated confidence interval procedure on efficiency with an empirical study on the drinking water in Morocco. The analysis was performed using some previous time effect functions in the literature and using some proposed sinusoidal time effect functions. Given the appeal of the numerical optimization due to the complexity of the log-likelihood functions in the presence of dependence and given the use of the parametric bootstrap in the construction of the confidence intervals, the computation was highly intense.

The study has demonstrated that the consideration of the dependence between the error components was strongly recommended given that the copula does not approach the product copula. It has revealed that the proposed time effect model which expresses that the amplitude of the time function decreases and that the time effect disappears over time is appropriate for our data according to the AIC criterion. The results revealed also a positive and significant time effect on the technical efficiency scores. It showed that the bias was important for several entities and that the efficiency scores were underestimated for thirty eight among the fifty provinces. It showed also a weak efficiency score and hence a weak bias corrected for almost all provinces, a very low national average of efficiency and therefore a weak performance of the ONEP's drinking

water. This might be due to several factors such as public management of the sector, the lack of cooperation projects to supply rural communities, the waste due to damaged pipes and the free distribution of drinking water through public fountains either in some rural districts or in informal areas of large cities. Unfortunately, the effect of factors on efficiency can not be measured in the absence of data in this regard. If data will be available, a study of the effect of the environmental variables will be done with the aim to identify the main determinants of the inefficiency in this domain in Morocco.

Will the total privatization of the sector yield to an improved efficiency? Some experiences as in Portugal indicate the poor performance of the private management in comparison with the public one as shown in De Witte and Marques (2008). Otherwise, it should be noted that business creation and economic development of rural municipalities enable rural citizens to have a considerable income and to fund their partial or total need of water and pipes. In addition, as an open issue, it will be interesting to perform a frontier analysis on the drinking water quality management in the country and to compare the provinces performance in this regard.

Chapter 5

Conclusions

The research of this thesis was dedicated to measuring efficiency of decision making units and making statistical inference about it. Hence, starting with the basic theory in this area which should respect some theoretical hypotheses, we investigated the problem of a possible dependence between the two random components in the stochastic frontier analysis. The classical independence hypothesis may not hold when it is needed and consequently we proposed appropriate procedures for estimation and statistical inference without the independence hypothesis. Smith (2008) has proposed an SFA model allowing for dependence, but beyond delivering the point estimate of technical efficiency, we proposed an interval estimate using a bootstrap procedure which is an extension of the algorithm developed in Simar and Wilson (2010) to the dependence case. Our procedures handled both cross-sectional and panel data.

We illustrated the statistical inference on technical efficiency through some Moroccan development areas, namely the local financing of the rural districts in both the oriental region and on a national level, and the drinking water in Morocco. Real data were considered, accompanied in some cases with simulated examples when they were needed. The first kind of data was handled in the second chapter using the nonparametric approaches DEA and FDH, while the

second kind of data was subject of the third chapter which treated the cross-sectional stochastic frontier analysis. The drinking water data set constituted the main experimental analysis of the panel stochastic frontier analysis of the fourth chapter.

As for the choice of the data sets analyzed in this thesis, we note that it was not easy to find complete databases for Morocco. Besides, I was constrained to collect most of the data by myself from archives and statistical yearbooks for several years. Moreover, analyses which require environmental variables were not made for lack or unavailability of data. The shortage of complete data in Morocco is due to the fact that political officials do not pay a lot of attention to numerical data. Even if our country began to pay interest to the numerical data for several years, its collection and its use remained constraint. It is only recently that numerical data began to expand in the analysis of the behaviour of entities in various sectors and in the economic planning. It is hence not obvious to find huge databases ready for the analysis. Consequently, it turns out essential and highly recommended to create a national observatory for statistical data which supplies detailed databases given that the statistics supplied by the Statistics Direction merely have a macroeconomic aspect.

Moreover, for some data sets as the local finances, even if several inputs exist, the curse of dimensionality which implies that working in smaller dimensions tends to provide better estimates of the frontier as pointed out in Daraio and Simar (2007), had incited to reduce the frontier estimation to the simplest case where one output and one or two inputs are considered. When it was needed, as in the empirical study on local financing, inputs are aggregated into a single input given that it is a meaningful variable. If this was not the case, Mouchard and Simar (2002) proposed a basic method to find an input which best summarizes the information provided by all inputs. As for the drinking water illustration, no aggregation was made given that two inputs are considered. We mentioned here again the lack of data which did not allow us to determine the causes of the inefficiency in this area.

The novelty of this research regarding the data is to handle the efficiency of the Moroccan local government and drinking water given that, to the best of our knowledge, no study has been made in this sense previously. However,

results in these areas should be carefully considered given the unavailability of certain factors and variables which allow us to have reliable conclusions.

The estimation method considered in the DEA approach was a simplex method where the linear programming problem was solved n times, once for each DMU. In both cross-section and panel data SFA, the estimation methods were the OLS, COLS and MLE. In addition, the log-likelihood function being highly nonlinear in all studied cases, the numerical optimization was adopted as mentioned previously using the Nelder-Mead algorithm. In this context, it may be possible in the future to explore cases, if they exist, of SFA models with dependent error terms for which an analytical form of the likelihood is available, and also a list of some of those models which lead to complicated forms that cannot be handled analytically, and hence numerical optimization is required. Several attempts were established but without leading to a model with simple analytical form.

We investigated statistical inference on the technical efficiency estimates provided by various approaches and mainly the parametric SFA. Two adapted parametric bootstrap procedures were proposed, one for the cross-section case and the other one for the panel data case. These procedures are a modified and an extended Algorithm#3 of Simar and Wilson (2010) to the copula case. The main contribution here is at the second stage of the algorithm where the error term components were drawn dependently using a copula and, of course, the other stages were adapted for the panel data case. Implementations were easy but unfortunately their computations under the dependence hypothesis were intense and required an enormous time to build confidence intervals.

Another contribution of this research is the proposal and comparison of a variety of deterministic time effects models expressing the evolution of inefficiency over time. In particular, we allowed not only for exponentially disappearing inefficiencies, but also for periodic, trigonometric type functions which may represent cycles in the evolution of inefficiencies. One of the periodic models was chosen as the best model describing our data in comparison with some known models in the literature.

About the empirical results, the first study related to the nonparametric DEA and FDH techniques in the case of the local financing of the eastern region

has revealed the existence of variable returns to scale (VRS) of our data and that among ninety-one only three districts for DEA and eight districts for FDH are efficient. It showed also that several financially autonomous municipalities are not efficient; hence their problem does not consist in the financing but in their management. The chapter has supplied an answer to the relationship between the efficiency scores and the population size of the districts using the truncated regression and the Kendall's tau which have indicated an inverse relation between them.

The second study provided a bootstrap procedure to estimate confidence intervals for technical efficiencies when the two error terms are dependent in the SFA technique. Various copula functions are considered and the model where the noise term has a normal distribution, the inefficiency term has a half-normal distribution and where the dependence between them was expressed by the Clayton copula was selected according to the highest likelihood value. The main finding is that the technical efficiency scores under dependence are lower than under independence while the rank remains the same. This would imply that assuming independence, efficiencies may be overestimated. In addition, scores are covered by their associated confidence intervals and generally the range of each interval is rather small. As for the relation between the efficiency and the population size, the same deduction as the nonparametric approach was established according to the Kendall's tau associated to the copula function.

As for the third study, it was an extension of the previous one to the panel data which is certainly richer in information than the cross-sectional case. Also a bootstrap procedure for the confidence intervals of technical efficiency was provided and illustrated using drinking water data regarding water production and its sale. Two inputs are considered which are the number of subscriptions and the total amount of sales, and one output which is the water production for a period of seven years. The Clayton copula is considered to express the dependence between the two error terms and a number of time-varying models were proposed. The panel stochastic frontier model where the time-varying inefficiency is expressed by the fact that the amplitude decreases over time and that the time effect disappears at the end of the period under study is chosen using the minimum AIC criterion. The result shows mainly that all

provinces are technically inefficient and that the time effect on the efficiency scores is positive and statistically significant. Here also a parametric percentile bootstrap procedure was provided to estimate confidence intervals of efficiencies.

On the other hand, it is useful to indicate that for all techniques adopted in this research the efficiencies are relative and their values depend on the best DMU in the dataset. So, if another DMU which dominates this latter is introduced, the efficiency scores can change and will tend towards a reduction. In addition, these scores depend on the chosen distribution, on the variables handled and may also be sensitive to the outliers which is not the case for our DEA analysis. Then, in order to validate our models, several tests were performed and analyses on the sensitivity to outliers and model identifiability were established. Furthermore, several extensions of the proposed models are possible to explain the behavior of the inputs and the output in the frontier analysis.

Moreover, there is an ambiguity in the SFA approach when the residuals of the OLS estimates are right-skewed (positive skewness), this might indicate that there is no inefficiency, so all DMUs are efficient, or that the model is misspecified. To overcome and avoid the problem, we have increased the sample size in Chapter 3 and we have investigated another model in a forthcoming work which is in progress where the inefficiency term in the panel frontier model is an extended exponential or an extended half-normal distributions having the “wrong” skewness. This work is an extension of Daniel et al (2011) and Hafner et al. (2015) to the panel data case for the time varying inefficiency term. Hafner et al. (2015) defines the extended exponential distribution as $f_{\tilde{\mu}}(u_i) = \frac{1}{\tilde{\mu}} \cdot \exp\left\{-\frac{|u_i|}{\tilde{\mu}}\right\}$ where $\tilde{\mu} = -\mu$; and the extended half-normal distribution as the usual half-normal $f_{\tilde{\mu}}(u_i) = \frac{2}{\tilde{\mu}\sqrt{\pi/2}} \cdot \phi\left(\frac{u_i}{\tilde{\mu}\sqrt{\pi/2}}\right)$ with $\tilde{\mu} = -\mu = \sqrt{2/\pi}\sigma_U$ and ϕ is the standard normal probability density function.

Indeed, we intend to carry out an analysis including three essential aspects in the presence of a positive skewness. For all of them, the panel theoretical model and simulated illustrations will be established and then they will be applied to concrete data if possible.

Thus, the first aspect of our future research consists in the proposal of the panel data version of the extended exponential and extended half-normal

distributions of Hafner et al. (2015) under the independence of the error terms hypothesis. Indeed, theoretical expressions of the error term densities and those of the technical efficiencies will be provided for both distributions. So, when $v_{it} \sim N(0, \sigma_V^2)$, u_i is e.g. the extended exponential distribution and $u_{it} = \eta(t)u_i$, the density of the error term ϵ_i and technical efficiency TE_{it} are expressed as (see the Appendix E for more details):

$$g(\epsilon_i) = \frac{\sigma_* \exp\left\{-\frac{1}{2}a_{*i}\right\}}{(2\pi)^{(T-1)/2} \sigma_V^T \tilde{\mu}} \Phi\left(-\frac{\mu_{*i}}{\sigma_*}\right) \quad (5.0.1)$$

where

$$\begin{aligned} \sigma_*^2 &= \frac{\sigma_V^2 \tilde{\mu}}{\tilde{\mu} \sum_t \eta^2(t)} = \frac{\sigma_V^2}{\sum_t \eta^2(t)}, \\ \mu_{*i} &= -\frac{\tilde{\mu} \sum_t \epsilon_{it} \eta(t) + 2\tilde{\mu}^2 \sum_t \eta^2(t) - \sigma_V^2}{\tilde{\mu} \sum_t \eta^2(t)}, \\ a_{*i} &= \frac{\sum_t \epsilon_{it}^2 + 4\tilde{\mu} \sum_t \epsilon_{it} \eta(t) + 4\tilde{\mu}^2 \sum_t \eta^2(t)}{\sigma_V^2} - \frac{\mu_{*i}^2}{\sigma_*^2}, \end{aligned}$$

and

$$TE_{it} = B_{it}/B_i = B_{it}/g(\epsilon_i), \quad (5.0.2)$$

where

$$\begin{aligned} B_{it} &= \frac{\sigma_* \exp\left\{-\frac{1}{2}a_{*it}\right\}}{(2\pi)^{(T-1)/2} \sigma_V^T \tilde{\mu}} \Phi\left(-\frac{\mu_{*it}}{\sigma_*}\right), \\ \mu_{*it} &= -\frac{\tilde{\mu} \sum_t \epsilon_{it} \eta(t) + 2\tilde{\mu}^2 \sum_t \eta^2(t) - \sigma_V^2 + \tilde{\mu} \sigma_V^2 \eta(t)}{\tilde{\mu} \sum_t \eta^2(t)}, \\ a_{*it} &= \frac{\sum_t \epsilon_{it}^2 + 4\tilde{\mu} \sum_t \epsilon_{it} \eta(t) + 4\tilde{\mu}^2 \sum_t \eta^2(t) + 4\sigma_V^2 \tilde{\mu} \eta(t)}{\sigma_V^2} - \frac{\mu_{*it}^2}{\sigma_*^2}. \end{aligned}$$

As an alternative to the first aspect, we plan in the second to answer the question: Could the consideration of the dependence hypothesis between the error term components reduce or eliminate the positive skewness problem for panel data characterized by this feature? Here again, the usual exponential and

half-normal distributions will be considered in order to compare this approach and the previous one using popular model selection techniques.

Furthermore, it also seems important to find in the future a way which reduces the intense calculation time which has required in our studies the simultaneous use of several computers especially during the inference procedure of the bootstrap methods in the cases of cross-sectional SFA and panel SFA as well as during the coverage estimation. So, we intend to propose an R package which realizes these computations and which can be publicly available.

Finally, at the applied level, this PhD research showed generally the inefficiency of the handled sectors. In the light of this result, the local councilors and the decision-makers of the drinking water sector are constrained to revise their management policies. At the theoretical level, new time-varying inefficiency models and two procedures were proposed to solve the problem of the independence between the term error components when this hypothesis is not filled. Without doubt, there are several paths for further research on this topic and we will try to explore some of them in a near future.

Appendix A

Some definitions and proofs for the DEA approach

In this appendix, we show how to determine dual from the primal problem for the DEA approach but let us at first clarify the envelopment and the multiplier DEA program notions.

In the literature, to avoid all confusion between primal and dual DEA programs, the DEA program with input efficiency θ and the weight vector λ is called the envelopment DEA program, while the DEA program with the weights (or prices) of inputs and outputs v_j and u_k respectively is called the multiplier DEA program.

Knowing that the dual of the dual program is the primal, we suggest to determine the envelopment model from the multiplier one. For further simplifications we will work with a CRS model.

We know that the analytical relation between a primal and a dual model is presented by the following expression where the primal is a maximization and the dual is a minimization problem:

$$\begin{array}{ll} \text{Max } c^t.x & \text{Min } b^t.y \\ \text{s.t. } \left\{ \begin{array}{l} A.x \leq b \\ x \geq 0 \end{array} \right. & \Rightarrow \text{s.t. } \left\{ \begin{array}{l} A.y \geq c \\ y \geq 0 \end{array} \right. \end{array} \quad (\text{A.0.1})$$

where x , y , b and c are vectors and A is a matrix. The strong duality theorem

states that if the primal has an optimal solution x^* , then the dual also has an optimal solution y^* such that $c^t x^* = b^t y^*$. When the CRS model is considered, the u^* is dropped from (2.2.3) and we have

$$\begin{aligned} & \text{Max} \sum_{k=1}^q u_k y_{k0} \\ & \text{s.t.} \begin{cases} \sum_{j=1}^p v_j x_{j0} \leq 1, & \leftarrow \theta \\ \sum_{k=1}^q u_k y_{ki} - \sum_{j=1}^p v_j x_{ji} \leq 0, & i = 1, \dots, n \quad \leftarrow \lambda_i \\ u_k, v_j \geq 0 \quad \forall k, j \end{cases} \quad (\text{A.0.2}) \end{aligned}$$

Let θ be the variable corresponding to the first constraint and λ be the variable corresponding to the other n inequality constraints in (A.0.2). Then (A.0.2) $\Rightarrow$

$$\left\{ \begin{array}{l} \text{Max } u_1 y_{10} + \dots + u_k y_{k0} + \dots + u_q y_{q0} + v_1 0 + \dots + v_j 0 + \dots + v_p 0 \\ \text{s.t.} \left\{ \begin{array}{l} \begin{pmatrix} 0 & \dots & 0 & \dots & 0 & x_{10} & \dots & x_{j0} & \dots & x_{p0} \\ y_{11} & \dots & y_{k1} & \dots & y_{q1} & -x_{11} & \dots & -x_{j1} & \dots & -x_{p1} \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ y_{1i} & \dots & y_{ki} & \dots & y_{qi} & -x_{1i} & \dots & -x_{ji} & \dots & -x_{pi} \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ y_{1n} & \dots & y_{kn} & \dots & y_{qn} & -x_{1n} & \dots & -x_{jn} & \dots & -x_{pn} \end{pmatrix} \cdot \begin{pmatrix} u_1 \\ \vdots \\ u_k \\ \vdots \\ u_q \\ v_1 \\ \vdots \\ v_j \\ \vdots \\ v_p \end{pmatrix} \leq \begin{pmatrix} 1 \\ 0 \\ \vdots \\ 0 \\ \vdots \\ 0 \end{pmatrix} \\ u_1, \dots, u_k, \dots, u_q, v_1, \dots, v_j, \dots, v_p \geq 0 \end{array} \right. \end{array} \right. \quad (\text{A.0.3})$$

$$\Rightarrow \left\{ \begin{array}{l} \text{Min } 1\theta + 0\lambda_1 + \cdots + 0\lambda_i + \cdots + 0\lambda_n \\ \text{s.t.} \left\{ \begin{array}{l} \begin{pmatrix} 0 & y_{11} & \cdots & y_{1i} & \cdots & y_{1n} \\ \vdots & \vdots & \cdots & \vdots & \cdots & \vdots \\ 0 & y_{k1} & \cdots & y_{ki} & \cdots & y_{kn} \\ \vdots & \vdots & \cdots & \vdots & \cdots & \vdots \\ 0 & y_{q1} & \cdots & y_{qi} & \cdots & y_{qn} \\ \vdots & \vdots & \cdots & \vdots & \cdots & \vdots \\ x_{10} & -x_{11} & \cdots & -x_{1i} & \cdots & -x_{1n} \\ \vdots & \vdots & \cdots & \vdots & \cdots & \vdots \\ x_{j0} & -x_{j1} & \cdots & -x_{ji} & \cdots & -x_{jn} \\ \vdots & \vdots & \cdots & \vdots & \cdots & \vdots \\ x_{p0} & -x_{p1} & \cdots & -x_{pi} & \cdots & -x_{pn} \end{pmatrix} \cdot \begin{pmatrix} \theta \\ \vdots \\ \lambda_1 \\ \vdots \\ \lambda_i \\ \vdots \\ \lambda_n \end{pmatrix} \geq \begin{pmatrix} y_{10} \\ \vdots \\ y_{k0} \\ \vdots \\ y_{q0} \\ 0 \\ \vdots \\ 0 \\ \vdots \\ 0 \end{pmatrix} \\ \theta \geq 0, \lambda_1, \cdots, \lambda_i, \cdots, \lambda_n \geq 0 \end{array} \right. \end{array} \right. \quad (\text{A.0.4})$$

$$\begin{array}{l} \text{Min } 1\theta + \sum_{i=1}^n 0\lambda_i \\ \text{s.t.} \begin{cases} 0\theta + \sum_{i=1}^n \lambda_i y_{ki} \geq y_{k0}, \quad k = 1, \cdots, q \\ \theta x_{j0} - \sum_{i=1}^n \lambda_i x_{ji} \geq 0, \quad j = 1, \cdots, p \\ \theta \geq 0, \lambda_i \geq 0 \quad \forall i \end{cases} \end{array}$$

$$\Rightarrow \begin{array}{l} \text{Min } 1\theta \\ \text{s.t.} \begin{cases} \sum_{i=1}^n \lambda_i y_{ki} \geq y_{k0}, \quad k = 1, \cdots, q \\ \theta x_{j0} - \sum_{i=1}^n \lambda_i x_{ji} \geq 0, \quad j = 1, \cdots, p \\ \theta \geq 0, \lambda_i \geq 0 \quad \forall i \end{cases} \end{array}$$

Finally we find the program given in (2.2.1) in the CRS case.

Appendix B

Some definitions and properties of copula functions

In this appendix we give some definitions and properties of copula functions, which can be used to model the dependence between random variables in a general way, and the associated Kendall's τ .

Definition *An n -dimensional copula is a distribution function defined on $[0, 1]^n$ with standard uniform marginal distributions.*

For general properties of copulas, we refer to Nelsen (1999). The following fundamental theorem has mainly motivated the widespread use of copulas.

Sklar's theorem (1959): *Given a multidimensional distribution function F which has $F_1, \dots, F_n$ as marginals, there exists a copula C of dimension n such that for all $a = (a_1, \dots, a_n) \in \mathbb{R}^n$:*

$$F(a_1, \dots, a_n) = C(F_1(a_1), \dots, F_n(a_n)),$$

where $F_i(a_i) = w_i$ for all $i = 1, \dots, n$ and $C(\vec{w}) = C(w_1, \dots, w_n)$ is a joint distribution with uniform marginals. Furthermore, if marginals are continuous, the copula C is unique and for all $w \in [0, 1]^n$ we can write

$$C(\vec{w}) = F(F_1^{-1}(w_1), \dots, F_n^{-1}(w_n)).$$

The function $\Pi(w_1, w_2) = w_1 \cdot w_2$ is called the product copula and has an im-

portant statistical interpretation: Let $w_1 = F_1(u)$ and $w_2 = F_2(v)$, $\Pi(w_1, w_2)$ is the copula of U and V if and only if they are independent.

Besides, the Kendall's tau, denoted τ , is a nonparametric statistical test to establish whether two random variables (X, Y) of n observations may be regarded as statistically dependent. It is hence a statistic used to measure the association between two variables using the rank correlation. For n pairs $(x_1, y_1), \dots, (x_n, y_n)$ and in the absence of ties, it is defined generally as the probability of concordance minus the probability of discordance expressed by

$$\frac{n_c}{\frac{1}{2}n(n-1)} - \frac{n_d}{\frac{1}{2}n(n-1)}, \quad (\text{B.0.1})$$

where n_c is the number of concordant pairs, n_d is the number of discordant pairs and $n_c + n_d = \binom{2}{n} = \frac{1}{2}n(n-1)$ is the total number of pairs. The pairs (x_i, y_i) and (x_j, y_j) are called concordant if $\text{sgn}(x_i - x_j) = \text{sgn}(y_i - y_j)$ and discordant if $\text{sgn}(x_i - x_j) = -\text{sgn}(y_i - y_j)$, where sgn means the sign function which is either positive or negative (it can not be zero because we have assumed the absence of ties). If ties are observed, pairs where at least one sign among the two compared signs is zero are neither concordant nor discordant and they are omitted from the total number of pairs to compute the Kendall's tau. Furthermore, $\tau \in [-1, 1]$ and a value of zero indicates the independence between X and Y .

In the following we give expressions for some bivariate copulas, their densities and their Kendall's τ . The copulas considered in this work are the Ali-Mikhail-Haq (AMH), Clayton, Fairlie-Gumbel-Morgenstern (FGM), Frank and Gaussian copulas. They all nest the product copula (i.e. independence) as a special case. For more details, see e.g. Genest and Favre (2007) and Nelsen (1999).

B.1 Gaussian copula

The Gaussian copula is given by

$$C_\theta(w_1, w_2) = \Phi_{2,\theta}(\Phi^{-1}(w_1), \Phi^{-1}(w_2)) \quad (\text{B.1.1})$$

and the Gaussian copula density function is obtained by differentiation w.r.t. w_1 and w_2 as

$$c_\theta(w_1, w_2) = \frac{\phi_{2,\theta}(\Phi^{-1}(w_1), \Phi^{-1}(w_2))}{\phi(\Phi^{-1}(w_1)) \cdot \phi(\Phi^{-1}(w_2))} = \frac{\phi_{2,\theta}(t_1, t_2)}{\phi(t_1) \cdot \phi(t_2)} \quad (\text{B.1.2})$$

with $t_i = \Phi^{-1}(w_i)$ for all $i = 1, 2$, and where $w_1 = F_1(u)$ and $w_2 = F_2(v)$ are the marginal distributions of u and v , respectively, θ is equal to the correlation coefficient, $\phi(t_i) = \Phi'(t_i)$ denotes the standard normal probability density function (p.d.f.), Φ the cumulative distribution function (c.d.f.) of the univariate standard normal distribution, and $\Phi_{2,\theta}$ denotes the c.d.f. of a bivariate Gaussian random variable with correlation θ and whose marginals are standard normal. The function $\Phi_{2,\theta}$ does not have a closed form expression, but it can be evaluated numerically. Furthermore, being in the class of elliptical distributions, the Gaussian copula is symmetric.

As described in Nelsen (1999) and Fredricks and Nelsen (2007), Kendall's τ as a function of the parameter θ for the Gaussian copula is expressed by $\tau = \frac{2}{\pi} \sin^{-1}(\theta)$. It verifies the relationship $-1 \leq \tau \leq 1$ and it indicates the independence when $\theta = 0$.

B.2 FGM copula

The Fairlie-Gumbel-Morgenstern copula, denoted FGM copula is the only copula which has a functional form as a second order polynomial in w_1 and w_2 . This FGM copula is defined in the bivariate case as

$$C_\theta(w_1, w_2) = w_1 w_2 P_\theta(w_1, w_2), \quad \theta \in [-1, 1] \quad (\text{B.2.1})$$

where the polynomial $P_\theta(w_1, w_2) = 1 + \theta(1 - w_1)(1 - w_2)$, hence

$$C_\theta(w_1, w_2) = w_1 w_2 [1 + \theta(1 - w_1)(1 - w_2)], \quad \theta \in [-1, 1]. \quad (\text{B.2.2})$$

Moreover, the copula density is given by

$$c_\theta(w_1, w_2) = \frac{\partial^2 C_\theta(w_1, w_2)}{\partial w_1 \partial w_2} = 1 + \theta - 2\theta w_1 - 2\theta w_2 + 4w_1 w_2, \quad (\text{B.2.3})$$

The functions C_θ and c_θ are respectively the c.d.f. and the p.d.f. of the FGM copula. The product copula is obtained as a special case for $\theta = 0$. The FGM copula is symmetric (exchangeable), meaning that $C_\theta(w_1, w_2) = C_\theta(w_2, w_1)$ for all $(w_1, w_2) \in I^2$ and that (w_1, w_2) and (w_2, w_1) are identically distributed. The associated Kendall's tau is $\tau = \frac{2}{9}\theta$ with $-\frac{2}{9} \leq \tau \leq \frac{2}{9}$.

B.3 Ali-Mikhail-Haq (AMH) copula

The AMH copula can represent both positive and negative dependence. The distribution and the density expressions of this copula are respectively

$$C_\theta(w_1, w_2) = \frac{w_1 w_2}{1 - \theta(1 - w_1)(1 - w_2)}, \quad \theta \in [-1, 1] \quad (\text{B.3.1})$$

$$c_\theta(w_1, w_2) = A_\theta(w_1, w_2) / B_\theta(w_1, w_2), \quad \theta \in [-1, 1] \quad (\text{B.3.2})$$

where

$$A_\theta(w_1, w_2) = -(1 - 2\theta + \theta^2 w_1 w_2 + \theta w_1 w_2 - \theta^2 w_2 + \theta^2 + \theta w_1 + \theta w_2 - \theta^2 w_1), \quad (\text{B.3.3})$$

$$B_\theta(w_1, w_2) = (-1 + \theta - \theta w_1 - \theta w_2 + \theta w_1 w_2)^3. \quad (\text{B.3.4})$$

If $\theta = 0$, then the two variables U and V are independent. For this copula $\tau = \frac{3\theta-2}{3\theta} - \frac{2(1-\theta)^2 \ln(1-\theta)}{3\theta^2}$ with $(5 - 8 \ln(2)) / 3 \leq \tau \leq 1/3$ for $\theta \neq 0$ and $\tau(0) = 0$ which corresponds to the product copula.

B.4 Clayton copula

The distribution function of the Clayton copula is defined by

$$C_\theta(w_1, w_2) = (w_1^{-\theta} + w_2^{-\theta} - 1)^{-1/\theta}, \quad \theta > 0 \quad (\text{B.4.1})$$

The expression of the copula density is

$$c_{\theta}(w_1, w_2) = w_1^{-1-\theta} w_2^{-1-\theta} \left((w_1^{-\theta} + w_2^{-\theta} - 1)^{-2-1/\theta} \right) (1 + \theta), \quad \theta > 0 \quad (\text{B.4.2})$$

As the parameter θ approaches zero, the two variables U and V become independent and the product copula is obtained as the limiting case:

$\lim_{\theta \rightarrow 0} C_{\theta}(w_1, w_2) = \Pi(w_1, w_2)$. As for the Kendall's tau, it is given by $0 < \tau = \frac{\theta}{\theta+2} < 1$.

B.5 Frank copula

The cdf of the Frank copula is given for all $\theta \in \mathbb{R}^*$ by

$$C_{\theta}(w_1, w_2) = -\frac{1}{\theta} \ln \left(1 + \frac{(\exp\{-\theta w_1\} - 1)(\exp\{-\theta w_2\} - 1)}{\exp\{-\theta\} - 1} \right), \quad (\text{B.5.1})$$

and the corresponding density by

$$c_{\theta}(w_1, w_2) = D_{\theta}(w_1, w_2) / E_{\theta}(w_1, w_2), \quad \theta \in \mathbb{R}^* \quad (\text{B.5.2})$$

where D and E are defined as

$$D_{\theta}(w_1, w_2) = \exp\{(1 + w_1 + w_2)\theta\} (\exp\{\theta\} - 1)\theta, \quad (\text{B.5.3})$$

$$E_{\theta}(w_1, w_2) = \left(\exp\{\theta\} + \exp\{(w_1 + w_2)\theta\} - \exp\{\theta + w_1\theta\} - \exp\{\theta + w_2\theta\} \right)^2. \quad (\text{B.5.4})$$

As for the Clayton copula, if $\theta \rightarrow 0$, then the product copula is obtained as the limiting case. Its Kendall's tau is defined as $\tau = 1 - \frac{4}{\theta} (1 - D_1(\theta))$ where $-1 \leq \tau \leq 1$ and $D_k(\theta) = \frac{k}{\theta^k} \int_0^{\theta} \frac{t^k}{e^t - 1} dt$ is the Debye function for any positive integer k .

Appendix C

MLE with independent error terms and time-varying technical efficiency

C.1 The time-constant efficiency

When the efficiency effects are assumed constant over time, the frontier model is called a fixed-effects model. It is assumed in this model that u_i are fixed but could be correlated with the regressors. The model constitutes the simplest panel data model and is written as

$$y_{it} = f(x_{it}, \beta) + v_{it} - u_i, \quad i = 1, \dots, n ; t = 1, \dots, T \quad (\text{C.1.1})$$

In the case where $y_{it} = \delta_i + \sum_{j=1}^l \beta_j x_{ij t} + v_{it}$, where $\delta_i = \beta_0 - u_i$, Schmidt and Sickles (1984) noted that in finite samples (small T) $\hat{\beta}_0$ is likely to be biased upward, which implies that efficiency is underestimated.

The time-constant model is a random-effects model when u_i are randomly distributed with a constant mean and a constant variance, and they are uncorrelated with v_{it} and with the regressors. The model is when the Cobb-Douglas

function is adopted

$$y_{it} = \beta_0^* + \sum_{j=1}^l \beta_j x_{ijt} + v_{it} - u_i^*, \quad i = 1, \dots, n; \quad t = 1, \dots, T, \quad (\text{C.1.2})$$

where $\beta_0^* = \beta_0 - E(u_i)$ and $u_i^* = u_i - E(u_i)$

C.2 Maximum Likelihood Estimation for the model with independent errors terms and time-varying technical efficiency

As described in Kumbhakar and Lovell (2000), for the model with time-varying technical efficiency given by

$$y_{it} = \beta_{0t} + \sum_{j=1}^l \beta_j x_{ijt} + v_{it} - u_{it} \quad (\text{C.2.1})$$

where $u_{it} = \gamma_t u_i$ and where $v_{it} \sim iidN(0, \sigma_V^2)$, $u_i \sim iidN^+(0, \sigma_U^2)$, v_{it} and u_{it} are uncorrelated and they are independent from the regressors, the MLE is as follows: Defining $\epsilon_{it} = v_{it} - u_{it} = v_{it} - \gamma_t u_i$ and $\epsilon_i = (\epsilon_{i1}, \dots, \epsilon_{iT})'$, with

$$\begin{aligned} f(\epsilon_i) &= \int_0^\infty f(\epsilon_i, u_i) du_i \\ &= \int_0^\infty \prod_t f(\epsilon_{it} + \gamma_t u_i) f(u_i) du_i \\ &= \frac{2}{(2\pi)^{(T+1)/2} \sigma_V^T \sigma_U} \int_0^\infty \exp \left\{ -\frac{1}{2} \left[\frac{\sum_t (\epsilon_{it} + \gamma_t u_i)^2}{\sigma_V^2} + \frac{u_i^2}{\sigma_U^2} \right] \right\} du_i \\ &= \frac{2\sigma_* \exp \left\{ -\frac{1}{2} a_{*i} \right\}}{(2\pi)^{T/2} \sigma_V^T \sigma_U} \int_0^\infty \frac{1}{\sqrt{2\pi} \sigma_*} \exp \left\{ -\frac{1}{2\sigma_*^2} (u_i - \mu_{*i})^2 \right\} du_i \quad (\text{C.2.2}) \end{aligned}$$

Where

$$\begin{aligned} \int_0^\infty \frac{1}{\sqrt{2\pi}\sigma_*} \exp\left\{-\frac{1}{2\sigma_*^2}(u_i - \mu_{*i})^2\right\} du_i &= 1 - \Phi\left(-\frac{\mu_{*i}}{\sigma_*}\right) \\ \mu_{*i} &= \frac{(\sum_t \gamma_t \epsilon_{it}) \sigma_V^2}{\sigma_V^2 + \sigma_U^2 \sum_t \gamma_t^2} \\ \sigma_*^2 &= \frac{\sigma_V^2 \sigma_U^2}{\sigma_V^2 + \sigma_U^2 \sum_t \gamma_t^2} \\ a_{*i} &= \frac{1}{\sigma_V^2} \left[\sum_t \epsilon_{it}^2 - \frac{\sigma_U^2 (\sum_t \gamma_t \epsilon_{it})^2}{\sigma_V^2 + \sigma_U^2 \sum_t \gamma_t^2} \right] \end{aligned}$$

The log-likelihood function is given by

$$\begin{aligned} l = \ln(L) &= n \ln \sigma_* - \frac{n}{2} \ln \sigma_*^2 - \frac{1}{2} \sum_i^n a_{*i} - \frac{nT}{2} \ln \sigma_V^2 - \frac{n}{2} \ln \sigma_U^2 \\ &\quad + \sum_i^n \ln \left[1 - \Phi\left(-\frac{\mu_{*i}}{\sigma_*}\right) \right], \end{aligned} \tag{C.2.3}$$

which can be maximized to obtain maximum likelihood estimators of β , γ_t , σ_U^2 and σ_V^2 .

From the derivation of the log-likelihood function it is easy to show that $u_i \mid \epsilon_i \sim N^+(\mu_{*i}, \sigma_*^2)$. An estimator for u_i can be obtained from the mean or the mode of $u_i \mid \epsilon_i$, which are given by

$$E(u_i \mid \epsilon_i) = \mu_{*i} + \sigma_* \left[\frac{\phi(-\mu_{*i}/\sigma_*)}{1 - \Phi(-\mu_{*i}/\sigma_*)} \right],$$

$$M(u_i \mid \epsilon_i) = \begin{cases} \mu_{*i}, & \text{if } \sum_t \gamma_t \epsilon_{it} \geq 0, \\ 0, & \text{otherwise} \end{cases}$$

Once u_i has been estimated, u_{it} can be estimated from $\hat{u}_{it} = \hat{u}_i \hat{\gamma}_t$, where $\hat{u}_i$ is either $E(u_i \mid \epsilon_i)$ or $M(u_i \mid \epsilon_i)$ and the $\hat{\gamma}_t$ are maximum likelihood estimators of γ_t , $t = 1, \dots, T$, subject to a normalization such as $\gamma_1 = 1$ or $\gamma_T = 1$. The

minimum squared error predictor of technical efficiency is

$$\begin{aligned} E(\exp\{-u_{it}\} \mid \epsilon_i) &= E(\exp\{-u_i\gamma_t\} \mid \epsilon_i) \\ &= \frac{1 - \Phi\left(\gamma_t\sigma_* - \frac{\mu_{*i}}{\sigma_*}\right)}{1 - \Phi\left(-\frac{\mu_{*i}}{\sigma_*}\right)} \exp\left\{-\gamma_t\mu_{*i} + \frac{1}{2}\gamma_t^2\sigma_*^2\right\} \end{aligned}$$

Appendix D

R tools used

Several R packages have been used in this thesis. We will refer to the most useful of them knowing that each one needs several other packages to be loaded. Furthermore, useful commands and estimated computing time will be provided in this appendix.

D.1 R packages

- **akima** version 0.5-4: The package gives linear or cubic spline interpolation for irregular gridded data. It can provide a list of points which smoothly interpolate given data points, similar to a curve drawn by hand. The used command is *aspline*.
- **boot**: **boot** is a basic R package which proposes functions and datasets for bootstrapping and the simplex method for linear programming problems. The main functions used are *boot* and *simplex*.
- **copula** version 0.9-7: It proposes, among others, methods for density, distribution, random number generation of bivariate and multivariate dependence measures for elliptical, Archimedean, extreme value and some more copula families. The main used functions are *amhCopula*, *claytonCopula*, *dcopula*, *fgmCopula*, *frankCopula*, *normalCopula* and *rcopula*.
- **fdrtool** version 1.2.6: The package contains, among others, density, dis-

tribution, random number generation for the half-normal distribution. Used commands are *phalfnorm*, *qhalfnorm* and *rhalfnorm*.

- **FEAR** version 1.1: It is used for computing nonparametric efficiency estimates, making inference, and testing hypotheses in frontier models. Commands are provided in Wilson (2008) for bootstrapping as well as computation of some new, robust estimators of efficiency, etc. The main used functions are *ap*, *boot.sw98*, *bootstrap.ci*, *dea*, *fdh* and *trunc.reg*.
- **frontier** version 0.997-12: The package serves to establish the maximum likelihood estimation of stochastic frontier production and cost functions. Two specifications are available: the error components specification with time-varying efficiencies (Battese and Coelli, 1992) and a model specification in which the firm effects are directly influenced by a number of variables (Battese and Coelli, 1995). The main used functions are *efficiencies*, *frontier* and *sfa*.
- **mda** version 0.4-2: The mda package is a tool of mixture and flexible discriminant analysis, multivariate adaptive regression splines. The used commands are *mars* and *predict*.
- **msm** version 1.0: It is a package for multi-state Markov and Hidden Markov Models in Continuous Time. It proposes commands to compute density, distribution function, quantile function and random generation for the truncated Normal distribution with mean equal to mean and standard deviation equal to standard deviation before truncation, and truncated on the interval [lower, upper]. Used commands are *ptnorm*, *qtnorm* and *rtnorm*.
- **stats4**: It is a basic R package which proposes a significant number of statistical functions. It includes the function *mle* which is a function to calculate negative log-likelihood used to estimate parameters by the method of maximum likelihood for parametric approach using some methods as the simulated annealing method (SANN) and the Nelder-Mead method.

D.2 Estimated computation time

Estimated computation time depends on a number of factors such as the kind of data (cross-sectional or panel data), the number of individuals n under study, the number m of simulations used to approximate numerically the integral expressions, fixed at 1000 for our computations, and on the computer speed. For e.g. a Processor intel Core i7-2600 CPU@3.40GHz $\times$ 8 characterized by a memory of 16.75 GB and a CPU speed of 3.40 GHz, the estimated time is given as follows.

- Cross sectional: For $n = 1298$, computation needs on average 40 minutes for each bootstrap replication, times B the number of the bootstrap replications. Thus, it requires 467 hours ($40mn \times 700$) which is approximately 19.5 days ($467h/24$) using constantly a single computer without incidents or bugs. This has required the use of several computers at the same time working independently to reduce the computing time.
- Panel data: For $n = 50$ and $T = 7$, computation requires on average at least 334 hours ($40mn \times 500$) which means approximately 14 days on one computer.

Appendix E

Positive skewness in the panel stochastic frontier analysis

As an extension of Daniel et al (2011) and Hafner et al. (2015) to the panel data case, we present in this appendix the efficiency estimate in the presence of positive skewness under the independence between the components of the error term. We provide hence the error term density and the technical efficiency expressions in both cases of the extended exponential and the extended half-normal distributions considering a time varying inefficiency term.

E.1 The extended exponential distribution

At this point, we develop an expression of the density of the error to define the likelihood function of the panel SFA model described by formulas (4.2.2) and $u_{it} = \eta(t) u_i$ when v_{it} is normal and u_i is an extended exponential distribution which is an exponential distribution mirrored at zero (i.e. $\mu < 0$) and then shifted to the right by 2 times its mean such that it has the same mean $|\mu|$ as the original exponential distribution. It is characterized by a negative skewness as opposed to the positive skewness of the exponential distribution.

E.1.1 The density of the error

Considering the extended exponential distribution presented in Daniel et al (2011) and Hafner et al. (2015) and knowing that $\mu < 0$, the model is given by

$$\epsilon = v - u + 2E(u) \quad (\text{E.1.1})$$

with $E(u) = \mu$. The term u is a negative exponential and, being negative, $|u_i| = -u_i$, and $u_{it} = \eta(t) u_i$ for panel data. Let us denote $\tilde{\mu} := |\mu|$. Then,

$$\begin{aligned} g(\epsilon_i) &= \int_{-\infty}^0 f(\epsilon_i, u_i) du_i \\ &= \int_{-\infty}^0 f(\epsilon_{i1}, \dots, \epsilon_{iT}, u_i) du_i \\ &= \int_{-\infty}^0 \prod_t f_2(\epsilon_{it} + \eta(t) u_i - 2\eta(t) E(u)) f_1(u_i) du_i \\ &= \int_{-\infty}^0 \frac{1}{(2\pi)^{T/2} \sigma_V^T} \exp \left\{ -\frac{1}{2} \left[\frac{\sum_t (\epsilon_{it} + \eta(t) u_i - 2\eta(t) E(u))^2}{\sigma_V^2} \right] \right\} \frac{1}{\tilde{\mu}} \exp \left\{ -\frac{|u_i|}{\tilde{\mu}} \right\} du_i \\ &= \frac{1}{(2\pi)^{T/2} \sigma_V^T \tilde{\mu}} \int_{-\infty}^0 \exp \left\{ -\frac{1}{2} \left[\frac{\sum_t (\epsilon_{it} + \eta(t) u_i + 2\eta(t) \tilde{\mu})^2}{\sigma_V^2} \right] - \frac{|u_i|}{\tilde{\mu}} \right\} du_i \\ &= \frac{1}{(2\pi)^{T/2} \sigma_V^T \tilde{\mu}} \int_{-\infty}^0 \exp \left\{ -\frac{1}{2} \left[\frac{\sum_t (\epsilon_{it}^2 + \eta^2(t) u_i^2 + 2\epsilon_{it}\eta(t) u_i + 4\tilde{\mu}\epsilon_{it}\eta(t) + 4\tilde{\mu}^2\eta^2(t) u_i + 4\tilde{\mu}^2\eta^2(t))}{\sigma_V^2} + \frac{2|u_i|}{\tilde{\mu}} \right] \right\} du_i \\ &= \frac{1}{(2\pi)^{T/2} \sigma_V^T \tilde{\mu}} \exp \left\{ -\frac{1}{2} \frac{\sum_t \epsilon_{it}^2 + 4\tilde{\mu} \sum_t \epsilon_{it}\eta(t) + 4\tilde{\mu}^2 \sum_t \eta^2(t)}{\sigma_V^2} \right\} \\ &\quad \int_{-\infty}^0 \exp \left\{ -\frac{1}{2} \left[\frac{u_i^2 \sum_t \eta^2(t) + 2u_i \sum_t \epsilon_{it}\eta(t) + 4u_i \tilde{\mu} \sum_t \eta^2(t)}{\sigma_V^2} + \frac{2|u_i|}{\tilde{\mu}} \right] \right\} du_i \\ &= \frac{1}{(2\pi)^{T/2} \sigma_V^T \tilde{\mu}} \exp \left\{ -\frac{1}{2} \frac{\sum_t \epsilon_{it}^2 + 4\tilde{\mu} \sum_t \epsilon_{it}\eta(t) + 4\tilde{\mu}^2 \sum_t \eta^2(t)}{\sigma_V^2} \right\} \\ &\quad \int_{-\infty}^0 \exp \left\{ -\frac{1}{2} \left[\frac{u_i^2 \tilde{\mu} \sum_t \eta^2(t) + 2u_i \tilde{\mu} \sum_t \epsilon_{it}\eta(t) + 4u_i \tilde{\mu}^2 \sum_t \eta^2(t) + 2\sigma_V^2(-u_i)}{\sigma_V^2 \tilde{\mu}} \right] \right\} du_i \\ &= \frac{1}{(2\pi)^{T/2} \sigma_V^T \tilde{\mu}} \exp \left\{ -\frac{1}{2} \frac{\sum_t \epsilon_{it}^2 + 4\tilde{\mu} \sum_t \epsilon_{it}\eta(t) + 4\tilde{\mu}^2 \sum_t \eta^2(t)}{\sigma_V^2} \right\} \\ &\quad \int_{-\infty}^0 \exp \left\{ -\frac{1}{2} \left[\frac{u_i^2 \tilde{\mu} \sum_t \eta^2(t) + 2(\tilde{\mu} \sum_t \epsilon_{it}\eta(t) + 2\tilde{\mu}^2 \sum_t \eta^2(t) - \sigma_V^2) u_i}{\sigma_V^2 \tilde{\mu}} \right] \right\} du_i \\ &= \frac{\sigma_*}{(2\pi)^{(T-1)/2} \sigma_V^T \tilde{\mu}} \exp \left\{ -\frac{1}{2} \frac{\sum_t \epsilon_{it}^2 + 4\tilde{\mu} \sum_t \epsilon_{it}\eta(t) + 4\tilde{\mu}^2 \sum_t \eta^2(t)}{\sigma_V^2} \right\} \\ &\quad \int_{-\infty}^0 \frac{1}{\sigma_* \sqrt{2\pi}} \exp \left\{ -\frac{1}{2} \left[\frac{1}{\sigma_*^2} u_i^2 - \frac{2}{\sigma_*^2} \mu_{*i} u_i \right] \right\} du_i \end{aligned}$$

$$\begin{aligned}
&= \frac{\sigma_*}{(2\pi)^{(T-1)/2} \sigma_V^T \tilde{\mu}} \exp \left\{ -\frac{1}{2} \frac{\sum_t \epsilon_{it}^2 + 4\tilde{\mu} \sum_t \epsilon_{it} \eta(t) + 4\tilde{\mu}^2 \sum_t \eta^2(t)}{\sigma_V^2} \right\} \\
&\quad \int_{-\infty}^0 \frac{1}{\sigma_* \sqrt{2\pi}} \exp \left\{ -\frac{1}{2} \left[\frac{1}{\sigma_*^2} u_i^2 - \frac{2}{\sigma_*^2} \mu_{*i} u_i + \frac{\mu_{*i}^2}{\sigma_*^2} - \frac{\mu_{*i}^2}{\sigma_*^2} \right] \right\} du_i \\
&= \frac{\sigma_*}{(2\pi)^{(T-1)/2} \sigma_V^T \tilde{\mu}} \exp \left\{ -\frac{1}{2} \frac{\sum_t \epsilon_{it}^2 + 4\tilde{\mu} \sum_t \epsilon_{it} \eta(t) + 4\tilde{\mu}^2 \sum_t \eta^2(t)}{\sigma_V^2} - \frac{\mu_{*i}^2}{\sigma_*^2} \right\} \\
&\quad \int_{-\infty}^0 \frac{1}{\sigma_* \sqrt{2\pi}} \exp \left\{ -\frac{1}{2} \left[\frac{1}{\sigma_*^2} (u_i - \mu_{*i})^2 \right] \right\} du_i.
\end{aligned}$$

This can be written as

$$g(\epsilon_i) = \frac{\sigma_* \exp \left\{ -\frac{1}{2} a_{*i} \right\}}{(2\pi)^{(T-1)/2} \sigma_V^T \tilde{\mu}} \Phi \left(-\frac{\mu_{*i}}{\sigma_*} \right) \quad (\text{E.1.2})$$

where

$$\sigma_*^2 = \frac{\sigma_V^2 \tilde{\mu}}{\tilde{\mu} \sum_t \eta^2(t)} = \frac{\sigma_V^2}{\sum_t \eta^2(t)}, \quad (\text{E.1.3})$$

$$\mu_{*i} = -\frac{\tilde{\mu} \sum_t \epsilon_{it} \eta(t) + 2\tilde{\mu}^2 \sum_t \eta^2(t) - \sigma_V^2}{\tilde{\mu} \sum_t \eta^2(t)}, \quad (\text{E.1.4})$$

$$a_{*i} = \frac{\sum_t \epsilon_{it}^2 + 4\tilde{\mu} \sum_t \epsilon_{it} \eta(t) + 4\tilde{\mu}^2 \sum_t \eta^2(t)}{\sigma_V^2} - \frac{\mu_{*i}^2}{\sigma_*^2} \quad (\text{E.1.5})$$

Therefore, assuming the independence across DMUs, the log-likelihood function can be written, for $\vartheta = (\tilde{\mu}, \sigma_V, \beta_0, \beta, \underline{\eta}_k)$ where $\underline{\eta}_k$ is a vector of k parameters in $\eta(t)$, as

$$\begin{aligned}
l(\vartheta) &= \log L(\vartheta) \\
&= \sum_{i=1}^n \log g_{\vartheta}(\epsilon_i), \quad (\text{E.1.6})
\end{aligned}$$

E.1.2 Technical Efficiency expression

Technical Efficiency is expressed as

$$\begin{aligned}
TE_{it} &= E \left(\exp \{ -u_{it} \} \mid \epsilon_i \right) \\
&= E \left(\exp \{ -\eta(t) u_i + 2\eta(t) E(u) \} \mid (\epsilon_{i1}, \dots, \epsilon_{it}, \dots, \epsilon_{iT}) \right)
\end{aligned}$$

$$\begin{aligned}
&= \int_{-\infty}^0 \exp \{ -\eta(t) u_i + 2\eta(t) E(u) \} f_1(u_i | (\epsilon_{i1}, \dots, \epsilon_{it}, \dots, \epsilon_{iT})) du_i \\
&= \int_{-\infty}^0 \exp \{ -\eta(t) u_i + 2\eta(t) E(u) \} \frac{f(\epsilon_{i1}, \dots, \epsilon_{it}, \dots, \epsilon_{iT}, u_i)}{g(\epsilon_i)} du_i \\
&= \frac{1}{g(\epsilon_i)} \int_{-\infty}^0 \exp \{ -\eta(t) u_i + 2\eta(t) E(u) \} f(\epsilon_{i1}, \dots, \epsilon_{it}, \dots, \epsilon_{iT}, u_i) du_i \\
&= \left(\frac{\sigma_* \exp \left\{ -\frac{1}{2} a_{*it} \right\}}{(2\pi)^{(T-1)/2} \sigma_V^T \tilde{\mu}} \Phi \left(-\frac{\mu_{*it}}{\sigma_*} \right) \right) / g(\epsilon_i), \tag{E.1.7}
\end{aligned}$$

where

$$\mu_{*it} = -\frac{\tilde{\mu} \sum_t \epsilon_{it} \eta(t) + 2\tilde{\mu}^2 \sum_t \eta^2(t) - \sigma_V^2 + \tilde{\mu} \sigma_V^2 \eta(t)}{\tilde{\mu} \sum_t \eta^2(t)}, \tag{E.1.8}$$

$$a_{*it} = \frac{\sum_t \epsilon_{it}^2 + 4\tilde{\mu} \sum_t \epsilon_{it} \eta(t) + 4\tilde{\mu}^2 \sum_t \eta^2(t) + 4\sigma_V^2 \tilde{\mu} \eta(t)}{\sigma_V^2} - \frac{\mu_{*it}^2}{\sigma_*^2}. \tag{E.1.9}$$

E.2 The extended half-normal distribution

Under the same assumptions as the exponential distribution and knowing that $\tilde{\mu} = -\mu = \sqrt{2/\pi} \sigma_U$, the error density and the technical efficiency expressions in the half-normal distribution are the same as the exponential one as in formulas (E.1.2) and (E.1.7) with the same a_{*i} and a_{*it} and with

$$\sigma_*^2 = \frac{\tilde{\mu}^2 \pi \sigma_V^2}{\tilde{\mu}^2 \pi \sum_t \eta^2(t) + 2\sigma_V^2}, \tag{E.2.1}$$

$$\mu_{*i} = -\frac{\tilde{\mu}^2 \pi \sum_t \epsilon_{it} \eta(t) + 2\tilde{\mu}^3 \pi \sum_t \eta^2(t)}{\tilde{\mu}^2 \pi \sum_t \eta^2(t) + 2\sigma_V^2}, \tag{E.2.2}$$

$$\mu_{*it} = -\frac{\tilde{\mu}^2 \pi \sum_t \epsilon_{it} \eta(t) + 2\tilde{\mu}^3 \pi \sum_t \eta^2(t) + \tilde{\mu}^2 \pi \sigma_V^2 \eta(t)}{\tilde{\mu}^2 \pi \sum_t \eta^2(t) + 2\sigma_V^2}. \tag{E.2.3}$$

Bibliography

- Abbott, M. and Ch. Doucouliagos (February 2003). The efficiency of Australian universities: a data envelopment analysis. *Economics of Education Review*, Vol. 22, Issue 1, pp. 89-97.
- Abbott, M. and Ch. Doucouliagos (2009). Competition and efficiency : overseas students and technical efficiency in Australian and New Zealand universities. *Education economics*, Vol. 17, Issue 1, pp. 31-57.
- Adams, R.M., A.N. Berger, and R.C. Sickles (1999). Semiparametric approaches to stochastic panel frontiers with applications in the banking industry. *Journal of Business and Economic Statistics*, Vol. 17, Issue 3, pp. 349-358.
- Afonso, A. and S. Fernandes (2008). Assessing and explaining the relative efficiency of local government. *Journal of Socio-Economics*, Vol. 37, Issue 5, pp. 1946-1979.
- Afriat, S. (1972). Efficiency estimation of production functions. *International Economic Review*, Vol. 13, Issue 3, pp. 568-598.
- Agahi, H., K. Zarafshani, and A.M. Behjat (2008). The Effect of Crop Insurance on Technical Efficiency of Wheat Farmers in Kermanshah Province: A Corrected Ordinary Least Square Approach. *Journal of Applied Sciences*, Vol. 8, Issue 5, pp. 891-894.
- Aigner, D.J. and S.F. Chu, (1968). On estimating the industry production function. *American Economic Review*, Vol. 58, pp. 826-839.

- Aigner, D.J., C.A. Knox Lovell, and P. Schmidt (1977). Formulation and estimation of stochastic frontier production function models. *Journal of Econometrics*, Vol. 6, Issue 1, pp. 21-37.
- Ambapour, S. (2001). Estimation des frontières de production et mesures de l'efficacité technique. Working paper, Bureau d'Application des Méthodes Statistiques et Informatiques (BAMSI).
- Aragon, Y., A. Daouia, and C. Thomas-Agnan (2005). Nonparametric frontier estimation: A conditional quantile-based approach. *Econometric Theory*, Vol. 21, Issue 2, pp. 358-389.
- Athanassopoulos, A. and K. Triantis (1998). Assessing Aggregate Cost Efficiency and the Related Policy Implications for Greek Local Municipalities. *INFOR*, Vol. 36, Issue 3, pp. 66-83.
- Balaguer-Coll, M.T., D. Prior, and E. Tortosa-Ausina (2007). On the determinants of local government performance: A two-stage nonparametric approach. *European Economic Review*, Vol. 51, Issue 2, pp. 425-451.
- Banker, R.D., A. Charnes, and W.W. Cooper (1984). Some models for estimating technical and scale inefficiencies in data envelopment analysis. *Management Science*, Vol. 30, Issue 9, pp. 1078-1092.
- Banker, R.D., W.W. Cooper, L.M. Seiford, R.M. Thrall, and J. Zhu (2004). Returns to scale in different DEA models. *European Journal of Operational Research*, Vol. 154, Issue 2, pp. 345-362.
- Battese, G.E. and T.J. Coelli (1988). Prediction of firm level technical efficiencies with a generalized frontier production function and panel data. *Journal of Econometrics*, Vol. 38, Issue 3, pp. 387-399.
- Battese, G.E. and T.J. Coelli (1992). Frontier Production Functions, Technical Efficiency and Panel Data: With Application to Paddy Farmers in India. *The Journal of Productivity Analysis*, Vol. 3, Issue 1, pp. 153-169.

- Battese, G.E. and T.J. Coelli (1995). A Model for Technical Inefficiency Effects in a Stochastic Frontier Production Function for Panel Data. *Empirical Economics*, Vol. 20, pp. 325-332.
- Battese, G.E., T.J. Coelli, and T.C. Colby (1989). Estimation of Frontier Production Functions and the Efficiencies of Indian Farms Using Panel Data From ICRIAT's Village Level Studies. *Journal of Quantitative Economics*, Vol. 5, Issue 2, pp. 327-348.
- Battese, G.E., A. Heshmati, and L. Hjalmarsson (2000). Efficiency of labour use in Swedish banking industry: a stochastic frontier approach. *Empirical Economics*, Vol. 25, Issue 4, pp. 623-640.
- Bélisle, C.J.P. (1992). Convergence Theorems for a Class of Simulated Annealing Algorithms on $\mathbb{R}^d$. *Journal of Applied Probability*, Vol. 29, Issue 4, pp. 885-895.
- Bhat, C. and N. Eluru (2009). A Copula-Based Approach to Accommodate Residential Self-Selection Effects in Travel Behavior Modeling. *Transportation Research, Part B*, Vol. 43, Issue 7, pp. 749-765.
- Borger, B. and K. Kerstens (1996). Radial and nonradial measures of technical efficiency: An empirical illustration for Belgian local governments using an FDH reference technology. *Journal of Productivity Analysis*, Vol. 7, Issue 1, pp. 41-62.
- Bowlin, W.F., A. Charnes, W.W. Cooper and H.D. Sherman (1985). Data envelopment analysis and regression approaches to efficiency estimation and evaluation. *Annals of Operations Research*, Vol. 2, Issues 1-4, pp. 113-138.
- Cazals, C., J. P. Florens, and L. Simar (2002). Nonparametric frontier estimation: A robust approach. *Journal of Econometrics*, Vol. 106, Issue 1, pp. 1-25.
- Charnes, A., W.W. Cooper, and E. Rhodes (1978). Measuring the efficiency of decision making units. *European Journal of Operational Research*, Vol. 2, Issue 6, pp. 429-444.

- Christensen, L.R., D.W. Jorgenson, and L.J. Lau (1971). Conjugate Duality and the Transcendental Logarithmic Production Function. *Econometrica*, Vol. 39, Issue 4, pp. 255-256.
- Coelli, T.J. (1995a). Estimators and Hypothesis Tests for a Stochastic frontier function: A Monte Carlo Analysis. *Journal of Productivity Analysis*, Vol. 6, Issue 3, pp. 247-268.
- Coelli, T.J. (1995b). Recent developments in frontier modelling and efficiency measurement. *Australian Journal of Agricultural Economics*, Vol. 39, Issue 3, pp. 219-245.
- Coelli, T.J. (1996). A Guide to FRONTIER, Version 4.1: A Computer Program for Stochastic Frontier Production and Cost Function Estimation. Centre for efficiency and Productivity Analysis, CEPA Working Paper 96/07, Department of Econometrics, University of New England.
- Coelli, T. and S. Perelman (1999). A comparison of parametric and nonparametric distance functions: With application to European railways. *European Journal of Operational Research*, Vol. 117, Issue 2, pp. 326-339.
- Cornwell, C., P. Schmidt, and R.C. Sickles (1990). Production Frontiers with Cross-Sectional and Time-Series Variation in Efficiency Levels. *Journal of Econometrics*, Vol. 46, Issues 1-2, pp. 185-200.
- Cubbin, J. and G. Tzanidakis (June 1998). Regression versus data envelopment analysis for efficiency measurement: an application to the England and Wales regulated water industry. *Utilities Policy*, Vol. 7, Issue 2, pp. 75-85.
- Daniel, B.C., Hafner, M.C., H. Manner, and L. Simar (2011). Asymmetries in Business Cycles and the Role of Oil Production. (ISBA Discussion Paper; 2011/32), 2011. 23 p. <http://hdl.handle.net/2078.1/92536>.
- Daouia, A. and L. Simar (2005). Robust nonparametric estimators of monotone boundaries. *Journal of Multivariate Analysis*, Vol. 96, Issue 2, pp. 311-331.

- Daouia, A. and L. Simar (2007). Nonparametric efficiency analysis: A multivariate conditional quantile approach. *Journal of Econometrics*, Vol. 140, Issue 2, pp. 375-400.
- Daraio, C. and L. Simar (2005). Introducing environmental variables in nonparametric frontier models: A probabilistic approach. *Journal of Productivity Analysis*, Vol. 24, Issue 1, pp. 93-121.
- Daraio, C. and L. Simar (2007). *Advanced Robust and Nonparametric Methods in Efficiency Analysis: Methodology and Applications*. Springer.
- De Witte, K. and R.C. Marques (2008). Designing incentives in local public utilities, an international comparison of the drinking water sector. *Social Science Research Network SSRN* 1084807.
- Debreu, G. (1951). The coefficient of resource utilization. *Econometrica*, Vol. 19, Issue 3, pp. 273-292.
- Deprins, D., L. Simar, and H. Tulkens (1984). Measuring labor-efficiency in post offices. In Marchand, M., Pestieau, P. and Tulkens, H. (Eds.). *The performance of public enterprises: concepts and measurement*, Amsterdam: North-Holland, pp. 243-267.
- Dolton, P., O.D. Marcenaro, and L. Navarro (2001). The Effective Use of Student Time: A Stochastic Frontier Production Function Case Study. Center for the Economics of Education, ISSN 2045-6557.
- Efron, B. (1982). The Jackknife, the Bootstrap, and Other Resampling Plans. *CBMS-NSF Regional Conference Series in Applied Mathematics*, #38. Philadelphia: SIAM.
- El Mehdi, R. and C.M. Hafner (2014a). Local Government Efficiency: The Case of Moroccan Municipalities. *African Development Review*, Vol. 26, Issue 1, pp. 88-101.
- El Mehdi, R. and C.M. Hafner (2014b). Inference in stochastic frontier analysis with dependent error terms. *Mathematics and Computers in Simulation (MATCOM)*, Vol. 102, Issue C, pp. 104-116.

- Erber, G. (2006). Benchmarking Efficiency of Telecommunication Industries in the US and Major European Countries: A Stochastic Possibility Frontiers Approach. DIW Berlin Discussion Paper, Issue 621.
- Fan, Y., Q. Li, and A. Weersink (1996). Semiparametric estimation of stochastic production frontier models. *Journal of Business and Economic Statistics*, Vol. 14, Issue 4, pp. 460-468.
- Färe, R. (1984). The Dual Measurement of efficiency. *Journal of Economics*, Vol. 44, Issue 3, pp. 283-288.
- Faria, R.C., G.S. Souza, and T.B. Moreira (2005). Public Versus Private Water Utilities: Empirical Evidence for Brazilian Companies. *Economics Bulletin*, Vol. 8, Issue 2, pp. 1-7.
- Farrell, M.J. (1957). The measurement of productive efficiency. *Journal of the Royal Statistical Society*, Vol. 120, Issue 3, pp. 253-281.
- Fredricks, G.A. and R.B. Nelsen (2007). On the relationship between Spearman's rho and Kendall's tau for pairs of continuous random variables. *Journal of Statistical Planning and Inference*, Vol. 137, pp. 2143-2150.
- Gallant, A. Ronald (1984). The Fourier Flexible Form. *American Journal of Agricultural Economics*, Vol. 66, Issue 2, pp. 204-208.
- Genest, C. and A.C. Favre (2007). Everything You Always Wanted to Know about Copula Modeling but Were Afraid to Ask. *Journal of Hydrologic Engineering*, Vol. 12, Issue 4, pp. 347-368.
- Gijbels, I., E. Mammen, B.U. Park, and L. Simar (1999). On Estimation of Monotone and Concave Frontier Functions. *Journal of the American Statistical Association*, Vol. 94, Issue 445, pp. 220-228.
- Greene, W. (1980a). Maximum likelihood estimation of econometric frontier function. *Journal of Econometrics*, Vol. 13, Issue 1, pp. 27-56.
- Greene, W. (1980b). On the estimation of a flexible frontier production. *Journal of Econometrics*, Vol. 13, Issue 1, pp. 101-115.

- Greene, W. (2005). *Efficiency of Public Spending in Developing Countries: A Stochastic Frontier Approach*. Stern School of Business, New York University.
- Greene, W. (2010). A stochastic frontier model with correction for sample selection, *Journal of Productivity Analysis*, Vol. 34, Issue 1, pp. 15-24.
- Hafner, M.C., H. Manner, and L. Simar (2015). The “wrong skewness” problem in stochastic frontier models: A new approach. *Econometric Reviews*, forthcoming.
- Hall, P. and L. Simar (June 2002). Estimating a Change point, Boundary or Frontier in the Presence of Observation Error. *Journal of the American Statistical Association*, Vol. 97, Issue 458, pp. 523-534.
- Haney, A.B. and M.G. Pollitt (December 2009). Efficiency analysis of energy networks: An international survey of regulators. *Energy Policy*, Vol. 37, Issue 12, pp. 5814-5830.
- Hattori, T., T. Jamasb, and M. Pollitt (2005). Electricity Distribution in the UK and Japan: A Comparative Efficiency Analysis 1985-1998. *The Energy Journal*, Vol. 26, Issue 2, pp. 23-47.
- Henderson, D.J. and L. Simar (September 2005). A Fully Nonparametric Stochastic Frontier Model for Panel Data. Discussion paper 0525, Institut de Statistique, UCL.
- von Hirschhausen, Ch., A. Cullmann, and A. Kappeler (2006). Efficiency analysis of German electricity distribution utilities - non-parametric and parametric tests. *Applied Economics*, Vol. 38, Issue 21, pp. 2553-2566.
- Holzer, M., J. Fry, E. Charbonneau, N. Riccucci, and E. Burnash (2009). *Literature Review and Analysis Related to Measurement of Local Government Efficiency*. Rutgers–Newark, SPAA (School of Public Affairs and Administration).
- Jacobs, R. (2001). Alternative Methods to Examine Hospital Efficiency: Data Envelopment Analysis and Stochastic Frontier Analysis. *Health Care Management Science*, Vol. 4, Issues 2, pp. 103-115.

Jondrow, J., C. A. Knox Lovell, I. S. Materov, and P. Schmidt (1982). On the estimation of technical inefficiency in the stochastic frontier production function model. *Journal of Econometrics*, Vol. 19, Issues 2-3, pp. 233-238.

Kim, S. and Y.H. Lee (2006). The productivity debate of East Asia revisited: a stochastic frontier approach. *Applied Economics, Taylor and Francis Journals*, Vol. 38, Issue 14, pp. 1697-1706.

Kneip, A., B.U. Park, and L. Simar (1998). A Note on the Convergence of Nonparametric DEA Estimators for Production Efficiency Scores. *Econometric Theory*, Vol. 14, Issue 6, pp. 783-793.

Kneip, A. and L. Simar (1996). A General Framework for Frontier Estimation with Panel Data. *Journal of Productivity Analysis*, Vol. 7, Issues 2-3, pp. 187-212.

Kneip, A., L. Simar, and I. Van Keilegom (2015). Frontier estimation in the presence of measurement error with unknown variance. *Journal of Econometrics*, Vol. 184, Issue 2, pp. 379-393.

Koopmans, T.C. (1951). An analysis of production as an efficient combination of activities. In: T.C. Koopmans (Ed.), *Activity Analysis of Production and Allocation*, Cowles Commission for Research in Economics. Monograph Issue 13, Wiley, New York

Kumbhakar, S.C. (1990). Production frontiers, panel data, and time-varying technical inefficiency. *Journal of Econometrics*, Vol. 46, Issues 1-2, pp. 201-211.

Kumbhakar, S.C. and C.A. Knox Lovell (2000). *Stochastic Frontier Analysis*, first ed., Cambridge University Press, United Kingdom.

Kumbhakar, S.C., B.U. Park, L. Simar, and E.G. Tsionas (2007). Nonparametric stochastic frontiers: a local maximum likelihood approach. *Journal of Econometrics*, Vol. 137, Issue 1, pp. 1-27.

- Kuosmanen, T. and M. Kortelainen (2012). Stochastic Non-Smooth Envelopment of Data: Semi-Parametric frontier estimation subject to shape constraints. *Journal of Productivity Analysis*, Vol. 38, Issue 1, pp. 11-28.
- Lee, Y.H. and P. Schmidt (1993). A Production Frontier Model with Flexible Temporal Variation in Technical Inefficiency. *The Measurement of Productive Efficiency: Techniques and Applications*. Edited by H. Fried, C.A.K. Lovell and S. Schmidt, Oxford University Press, pp. 237-255.
- Linna, M. (1998). Measuring Hospital Cost Efficiency with Panel Data Models. *Health Economics*, Vol. 7, Issue 5, pp. 415-427.
- Linna, M. and U. Häkkinen (1998). Determinants of cost efficiency of Finnish hospitals: A comparison of DEA and SFA. Helsinki University of Technology, Systems Analysis Laboratory, Research Report A78 (1998).
- Linna, M., U. Häkkinen, M. Peltola, J. Magnussen, K. S. Anthun, S. Kittelsen, A. Roed, K. Olsen, E. Medin, and C. Rehnberg (December 2010). Measuring cost efficiency in the Nordic Hospitals - a cross-sectional comparison of public hospitals in 2002. *Health Care Management Science*, Vol. 13, Issue 4, pp. 346-357.
- Loikkanen, H. and I. Susiloto (2005). Cost Efficiency of Finnish Municipalities in Basic Service Provision 1994-2002. *Urban Public Economics Review*, Vol. 4, pp. 39-64.
- Lovell, C.A.K. (2006). Frontier analysis in healthcare. *International Journal of Healthcare Technology and Management*, Vol. 7, Issues 1-2, pp. 5-14.
- Mbangala, M. and S. Perelman (1997). L'efficacité technique des chemins de fer en Afrique Subsaharienne: une comparaison internationale par la méthode DEA. *Revue d'Economie du Développement*, Vol. 5, Issue 3, pp. 91-115.
- McAllister, P.H. and D. McManus (1993). Resolving the scale efficiency puzzle in banking. *Journal of Banking and Finance*, Vol. 17, Issues 2-3, pp. 389-405.

- Meeusen, W. and J. Van Den Broeck (1977). Efficiency Estimation From Cobb-Douglas Production Function With Composed Error. *International Economic Review*, Vol. 18, Issue 2, pp. 435-444.
- Mlambo, K. and M. Ncube (2011). Competition and Efficiency in the Banking Sector in South Africa. *African Development Review*, Vol. 23, Issue 4, pp. 4-15.
- Moore, A., J. Nolan, and G.F. Segal (2005). Putting Out the Trash: Measuring Municipal Service Efficiency in U.S. Cities. *Urban Affairs Review*, Vol. 41, Issue 2, pp. 237-259.
- Mouchard, M. and L. Simar (2002). Efficiency Analysis of Air Controllers: First Insights. Consulting report No. 0202, Institut de Statistique, UCL, Belgium.
- Nelder, J.A. and R. Mead (1965). A simplex method for function minimization. *Computer Journal*, Vol. 7, Issue 4, pp. 308-313.
- Nelsen, R. B. (1999). *An Introduction to Copulas*, first ed., Springer, New York.
- Ogundari, K., T.T. Amos, and V.O. Okoruwa (2012). A Review of Nigerian Agricultural Efficiency Literature, 1999-2011: What Does One Learn from Frontier Studies?, *African Development Review*, Vol. 24, Issue 1, pp. 93-106.
- Park, B.U., R. Sickles, and L. Simar (1998). Stochastic Panel Frontiers: A Semiparametric Approach. *Journal of Econometrics*, Vol. 84, Issue 2, pp. 273-301.
- Park, B.U. and L. Simar (September 1994). Efficient Semiparametric Estimation in a Stochastic Frontier Model. *Journal of the American Statistical Association*, Vol. 89, Issue 427, pp. 929-936.
- Park, B.U., L. Simar, and Ch. Weiner (2000). The FDH Estimator for Productivity Efficiency Scores: Asymptotic Properties. *Econometric Theory*, Vol. 16, Issue 6, pp. 855-877.
- Perelman, S. (1996). La mesure de l'efficacité des services publics. *Revue Française des Finances Publiques*, Issue 55, pp. 65-79.

- Prieto, A. and J. Zofio (2001). Evaluating Effectiveness in Public Provision of Infrastructure and Equipment: The Case of Spanish Municipalities. *Journal of Productivity Analysis*, Vol. 15, Issue 1, pp. 41-58.
- R Development Core Team (2011). R: A language and environment for statistical computing. R Foundation for Statistical Computing, Vienna, Austria. <http://www.r-project.org/>.
- Ritter, C. and L. Simar (1997). Pitfalls of normal-gamma stochastic frontier models. *Journal of Productivity Analysis*, Vol. 8, Issue 2, pp. 167-182.
- Rosen, M.A., M.N. Le, and I. Dincer (2005). Efficiency analysis of a cogeneration and district energy system. *Applied Thermal Engineering*, Vol. 25, Issue 1, pp. 147-159.
- Sampaio, A., C. Barros, and J. Ramajo (2005). Technical Inefficiency in Municipal Water Distribution Service: A Case Study for Portugal. *Anales de Economía Aplicada*, XIX Reunión Anual. Edições Asepelt (Associação de Economia Aplicada) Espanha Badajoz.
- Schmidt, P. and R.C. Sickles (1984). Production Frontiers and Panel Data. *Journal of Business & Economic Statistics*, Vol. 2, Issue 4, pp. 367-374.
- Shephard, R.W. (1970). Theory of Cost and Production Functions. Princeton University Press, Princeton, NJ.
- Simar, L. (1992). Estimating Efficiencies from Frontier Models with Panel Data: A Comparison of Parametric, Non-Parametric and Semi-Parametric Methods with Bootstrapping. *The Journal of Productivity Analysis*, Vol. 3, Issues 1-2, pp. 171-203.
- Simar, L. (2003a). Detecting outliers in frontier models: a simple approach, *Journal of Productivity Analysis*, Vol. 20, Issue 22, pp. 391-424.
- Simar, L. (2003b). How to Improve the Performances of DEA/FDH Estimators in the Presence of Noise?. Discussion papers of interdisciplinary research project 373, No. 2003,33.

- Simar, L. and A. Vanhems (2012). Probabilistic characterization of directional distances and their robust versions. *Journal of Econometrics*, Vol. 166, Issue 2, pp. 342-354.
- Simar, L., A. Vanhems, and P. W. Wilson (2012). Statistical inference for dea estimators of directional distances. *European Journal of Operational Research*, Vol. 220, Issue 3, pp. 853-864.
- Simar, L. and P.W. Wilson (1998). Sensitivity Analysis of Efficiency Scores: How to Bootstrap in Nonparametric Frontier Models. *Management Science*, Vol. 44, Issue 1, pp. 49-61.
- Simar, L. and P.W. Wilson (2000a). Statistical inference in nonparametric frontier models: the state of the art. *Journal of Productivity Analysis*, Vol. 13, Issue 1, pp. 49-78.
- Simar, L. and P.W. Wilson (2000b). A general methodology for bootstrapping in non-parametric frontier models. *Journal of Applied Statistics*, Vol. 27, Issue 6, pp. 779-802.
- Simar, L. and P.W. Wilson (2002). Nonparametric tests of returns to scale. *European Journal of Operational Research*, Vol. 139, Issue 1, pp. 115-132.
- Simar, L. and P.W. Wilson (2007). Estimation and inference in two-stage, semi-parametric models of production processes. *Journal of Econometrics*, Vol. 136, Issue 1, pp. 31-64.
- Simar, L. and P.W. Wilson (2010). Inferences from Cross-Sectional, Stochastic Frontier Models. *Econometric Reviews*, Vol. 29, Issue 1, pp. 62-98.
- Simar, L. and P.W. Wilson (2013). Estimation and Inference in Nonparametric Frontier Models: Recent Developments and Perspectives. *Foundations and Trends®Econometrics*, Vol. 5, Issues 3-4, pp. 183-337.
- Simar, L. and P.W. Wilson (2015). Statistical Approaches for Non-parametric Frontier Models: A Guided Tour. *International Statistical Review*, Vol. 83, Issue 1, pp. 77-110.

- Smith, M.D. (2008). Stochastic frontier models with dependent error components. *The Econometrics Journal*, Vol. 11, Issue 1, pp. 172-192.
- Stevens, P.A. (2005). A Stochastic Frontier Analysis of English and Welsh Universities. *Education Economics*, Vol. 13, Issue 4, pp. 355-374.
- Sueyoshi, T. (1999). DEA Duality on Returns to Scale (RTS) in Production and Cost Analyses: An Occurrence of Multiple Solutions and Differences Between Production-Based and Cost-Based RTS Estimates. *Management Science*, Vol. 45, Issue 11, pp. 1593-1608.
- Thanassoulis, E. (1993). A Comparison of Regression Analysis and Data Envelopment Analysis as Alternative Methods for Performance Assessments. *The Journal of the Operational Research Society*, Vol. 44, Issue 11, pp. 1129-1144.
- Thanassoulis, E. (2001). *Introduction to the theory and application of data envelopment analysis: A foundation text with integrated software*. Springer Verlag, 281 p.
- Thompson, R.G., L.N. Langemeier, C. Lee, E. Lee, and R.M. Thrall (October-November 1990). The role of multiplier bounds in efficiency analysis with application to Kansas farming. *Journal of Econometrics*, Vol. 46, Issues 1-2, pp. 93-108.
- Tran, N.A., G. Shively, and P. Preckel (2010). A new method for detecting outliers in Data Envelopment Analysis. *Applied Economics Letters*, Vol. 17, Issue 4, pp. 313-316.
- Tran, K.C. and E.G. Tsionas (September 2009). Estimation of nonparametric inefficiency effects stochastic frontier models with an application to British manufacturing. *Economic Modelling*, Elsevier, Vol. 26, Issue 5, pp. 904-909.
- Tupper, H.C. and M. Resende (2004). Efficiency and regulatory issues in the Brazilian water and sewage sector: an empirical study. *Utilities Policy*, Vol. 12, Issue 1, pp. 29-40.
- Vishwakarma, A. and M. Kulshrestha (2010). Stochastic Production Frontier Analysis of Water Supply Utility of Urban Cities in the State of Madhya

- Pradesh, India. *International journal of environmental sciences*, Vol. 1, Issue 3, pp. 357-367.
- Wheelock, C. D. and P. W. Wilson (2008). Non-parametric, unconditional quantile estimation for efficiency analysis with an application to Federal Reserve check processing operations. *Journal of Econometrics*, Vol. 145, Issue 1-2, pp. 209-225.
- Wilson, P.W. (1993). Detecting outliers in deterministic nonparametric frontier models with multiple outputs. *Journal of Business and Economic Statistics*, Vol. 11, Issue 3, pp. 319-323.
- Wilson, P.W. (2008). FEAR 1.0: A Software Package for Frontier Efficiency Analysis with R. *Socio-Economic Planning Sciences*, Vol. 42, pp. 247-254.
- Wilson, W. P. (2011). Asymptotic properties of some non-parametric hyperbolic efficiency estimators. In: I. Van Keilegom and P. W. Wilson (eds.): *Exploring Research Frontiers in Contemporary Statistics and Econometrics*. Berlin: Springer-Verlag, pp. 115-150.
- Worthington, A. (2000). Cost efficiency in Australian local government: A comparative analysis of mathematical programming and econometric approaches. *Financial Accountability and Management*, Vol. 16, Issue 3, pp. 201-224.
- Worthington, A.C. and B.E. Dollery (2002). Incorporating Contextual Information in Public Sector Efficiency Analyses: A Comparative Study of NSW Local Government. *Applied Economics*, Vol. 34, Issue 4, pp. 453-464.