

Modeling Data Irregularities and Structural Complexities in Data Envelopment Analysis

Modeling Data Irregularities and Structural Complexities in Data Envelopment Analysis

Edited by

Joe Zhu

Worcester Polytechnic Institute, U.S.A.

Wade D. Cook

York University, Canada



Springer

Joe Zhu
Worcester Polytechnic Institute
Worcester, MA, USA

Wade D. Cook
York University
Toronto, ON, Canada

Library of Congress Control Number: 2007925039

ISBN 978-0-387-71606-0

e-ISBN 978-0-387-71607-7

Printed on acid-free paper.

© 2007 by Springer Science+Business Media, LLC

All rights reserved. This work may not be translated or copied in whole or in part without the written permission of the publisher (Springer Science+Business Media, LLC, 233 Spring Street, New York, NY 10013, USA), except for brief excerpts in connection with reviews or scholarly analysis. Use in connection with any form of information storage and retrieval, electronic adaptation, computer software, or by similar or dissimilar methodology now known or hereafter developed is forbidden.

The use in this publication of trade names, trademarks, service marks and similar terms, even if they are not identified as such, is not to be taken as an expression of opinion as to whether or not they are subject to proprietary rights.

9 8 7 6 5 4 3 2 1

springer.com

To *Alec and Marsha Rose*

CONTENTS

1	Data Irregularities and Structural Complexities in DEA	1
	Wade D. Cook and Joe Zhu	
2	Rank Order Data in DEA	13
	Wade D. Cook and Joe Zhu	
3	Interval and Ordinal Data	35
	Yao Chen and Joe Zhu	
4	Variables with Negative Values in DEA	63
	Jesús T. Pastor and José L. Ruiz	
5	Non-Discretionary Inputs	85
	John Ruggiero	
6	DEA with Undesirable Factors	103
	Zhongsheng Hua and Yiwen Bian	
7	European Nitrate Pollution Regulation and French Pig Farms' Performance	123
	Isabelle Piot-Lepetit and Monique Le Moing	
8	PCA-DEA	139
	Nicole Adler and Boaz Golany	
9	Mining Nonparametric Frontiers	155
	José H. Dulá	

10	DEA Presented Graphically Using Multi-Dimensional Scaling	171
	Nicole Adler, Adi Raveh and Ekaterina Yazhemsky	
11	DEA Models for Supply Chain or Multi-Stage Structure	189
	Wade D. Cook, Liang Liang, Feng Yang, and Joe Zhu	
12	Network DEA	209
	Rolf Färe, Shawna Grosskopf and Gerald Whittaker	
13	Context-Dependent Data Envelopment Analysis and its Use	241
	Hiroshi Morita and Joe Zhu	
14	Flexible Measures—Classifying Inputs and Outputs	261
	Wade D. Cook and Joe Zhu	
15	Integer DEA Models	271
	Sebastián Lozano and Gabriel Villa	
16	Data Envelopment Analysis with Missing Data	291
	Chiang Kao and Shiang-Tai Liu	
17	Preparing Your Data for DEA	305
	Joe Sarkis	
	About the Authors	321
	Index	331

Chapter 1

DATA IRREGULARITIES AND STRUCTURAL COMPLEXITIES IN DEA

Wade D. Cook¹ and Joe Zhu²

¹*Schulich School of Business, York University, Toronto, Ontario, Canada, M3J 1P3, wcook@shulich.yorku.ca*

²*Department of Management, Worcester Polytechnic Institute, Worcester, MA 01609, jzhu@wpi.edu*

Abstract: Over the recent years, we have seen a notable increase in interest in data envelopment analysis (DEA) techniques and applications. Basic and advanced DEA models and techniques have been well documented in the DEA literature. This edited volume addresses how to deal with DEA implementation difficulties involving data irregularities and DMU structural complexities. Chapters in this volumes address issues including the treatment of ordinal data, interval data, negative data and undesirable data, data mining and dimensionality reduction, network and supply chain structures, modeling non-discretionary variables and flexible measures, context-dependent performance, and graphical representation of DEA.

Key words: Data Envelopment Analysis (DEA), Ordinal Data, Interval Data, Data Mining, Efficiency, Flexible, Supply Chain, Network, Undesirable

1. INTRODUCTION

Data envelopment analysis (DEA) was introduced by Charnes, Cooper and Rhodes (CCR) in 1978. DEA measures the relative efficiency of peer decision making units (DMUs) that have multiple inputs and outputs, and has been applied in a wide range of applications over the past 25 years, in settings that include hospitals, banks, maintenance crews, etc.; see Cooper, Seiford and Zhu (2004).

As DEA attracts ever-growing attention from practitioners, its application and use become a very important issues. It is, therefore, important to deal with computation/data issues in DEA. These include, for example, how to deal with inaccurate data, qualitative data, outliers, undesirable factors, and many others. It is as well critical, from a managerial perspective, to be able to visualize DEA results, when the data are more than 3-dimensional.

The current volume presents a collection of articles that address data issues in the application of DEA, and special problem structures with respect to the nature of DMUs.

2. DEA MODELS

In this section, we present some basic DEA models that will be used in later chapters. For a more detailed discussion on these and other DEA models, the reader is referred to Cooper, Seiford and Zhu (2004), and other DEA textbooks.

Suppose we have a set of n peer DMUs, $\{DMU_j : j = 1, 2, \dots, n\}$, which produce multiple outputs y_{rj} , ($r = 1, 2, \dots, s$), by utilizing multiple inputs x_{ij} , ($i = 1, 2, \dots, m$). When a DMU_o is under evaluation by the CCR ratio model, we have (Charnes, Cooper and Rhodes, 1978)

$$\begin{aligned}
 \max \quad & \frac{\sum_{r=1}^s \mu_r y_{ro}}{\sum_{i=1}^m v_i x_{io}} \\
 \text{s.t.} \quad & \frac{\sum_{r=1}^s \mu_r y_{rj}}{\sum_{i=1}^m v_i x_{ij}} \leq 1, \quad j = 1, 2, \dots, n \\
 & \mu_r, v_i \geq 0, \quad \forall r, i
 \end{aligned} \tag{1}$$

In this model, inputs x_{ij} and outputs y_{rj} are observed non-negative data¹, and μ_r and v_i are the unknown weights, or decision variables.

A fully rigorous development would replace $\mu_r, v_i \geq 0$ with

¹ For the treatment of negative input/output data, please see Chapter 4.

$$\frac{\mu_r}{\sum_{i=1}^m v_i x_{io}}, \frac{\mu_r}{\sum_{i=1}^m v_i x_{io}} \geq \varepsilon > 0$$

where ε is a non-Archimedean element smaller than any positive real number.

Model (1) can be converted into a linear programming problem

$$\begin{aligned} & \max \sum_{r=1}^s \mu_r y_{ro} \\ & \text{subject to} \\ & \sum_{r=1}^s \mu_r y_{rj} - \sum_{i=1}^m v_i x_{ij} \leq 0, \text{ all } j \\ & \sum_{i=1}^m v_i x_{io} = 1 \\ & \mu_r, v_i \geq 0 \end{aligned} \tag{2}$$

In this model, the weights are usually referred to as multipliers. Therefore, model (2) is also called a multiplier DEA model. The dual program to (2) can be expressed as

$$\begin{aligned} & \theta^* = \min \theta \\ & \text{subject to} \\ & \sum_{j=1}^n x_{ij} \lambda_j \leq \theta x_{io} \quad i = 1, 2, \dots, m; \\ & \sum_{j=1}^n y_{rj} \lambda_j \geq y_{ro} \quad r = 1, 2, \dots, s; \\ & \lambda_j \geq 0 \quad j = 1, 2, \dots, n. \end{aligned} \tag{3}$$

Model (3) is referred to as the envelopment model. To illustrate the concept of envelopment, we consider a simple numerical example used in Zhu (2003) as shown in Table 1-1 where we have five DMUs representing five supply chain operations. Within a week, each DMU generates the same profit of \$2,000 with a different combination of supply chain cost and response time.

Table 1-1. Supply Chain Operations Within a Week

DMU	Inputs		Output
	Cost (\$100)	Response time (days)	Profit (\$1,000)
1	1	5	2
2	2	2	2
3	4	1	2
4	6	1	2
5	4	4	2

Source: Zhu (2003).

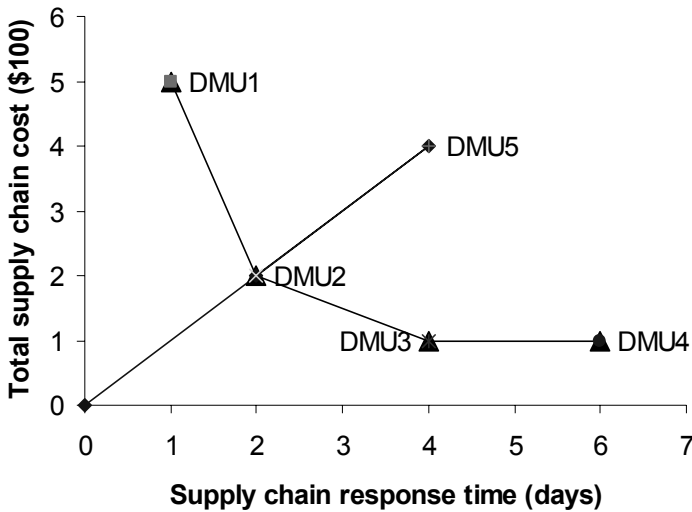


Figure 1-1. Five Supply Chain Operations

Figure 1-1 presents the five DMUs and the piecewise linear DEA frontier. DMUs 1, 2, 3, and 4 are on the frontier--or the *envelopment* frontier. If we apply model (3) to DMU5, we have,

$$\begin{aligned}
 &\text{Min } \theta \\
 &\text{Subject to} \\
 &1 \lambda_1 + 2\lambda_2 + 4\lambda_3 + 6\lambda_4 + 4\lambda_5 \leq 4\theta \\
 &5 \lambda_1 + 2\lambda_2 + 1\lambda_3 + 1\lambda_4 + 4\lambda_5 \leq 4\theta \\
 &2 \lambda_1 + 2\lambda_2 + 2\lambda_3 + 2\lambda_4 + 2\lambda_5 \geq 2 \\
 &\lambda_1, \lambda_2, \lambda_3, \lambda_4, \lambda_5 \geq 0
 \end{aligned}$$

This model has the unique optimal solution of $\theta^* = 0.5$, $\lambda_2^* = 1$, and $\lambda_j^* = 0$ ($j \neq 2$), indicating that DMU5 needs to reduce its cost and response time to

the amounts used by DMU2 if it is to be efficient This example indicates that technical efficiency for DMU5 is achieved at DMU2.

Now, if we apply model (3) to DMU4, we obtain $\theta^* = 1$, $\lambda_4^* = 1$, and $\lambda_j^* = 0$ ($j \neq 4$), indicating that DMU4 is on the frontier. However, Figure 1-1 indicates that DMU4 can still reduce its response time by 2 days to achieve coincidence with DMU3. This input reduction is usually called input slack.

The nonzero slack can be found by using the following model

$$\begin{aligned} & \max \sum_{i=1}^m s_i^- + \sum_{r=1}^s s_r^+ \\ & \text{subject to} \\ & \sum_{j=1}^n x_{ij} \lambda_j + s_i^- = \theta^* x_{io} \quad i = 1, 2, \dots, m; \\ & \sum_{j=1}^n y_{rj} \lambda_j - s_r^+ = y_{ro} \quad r = 1, 2, \dots, s; \\ & \lambda_j, s_i^-, s_r^+ \geq 0 \quad \forall i, j, r \end{aligned} \quad (4)$$

where θ^* is determined by model (3) and is fixed in model (4).

For DMU4 with $\theta^* = 1$, model (4) yields the following model,

$$\begin{aligned} & \text{Max } s_1^- + s_2^- + s_1^+ \\ & \text{Subject to} \\ & 1 \lambda_1 + 2\lambda_2 + 4\lambda_3 + 6\lambda_4 + 4\lambda_5 + s_1^- = 6 \theta^* = 6 \\ & 5 \lambda_1 + 2\lambda_2 + 1\lambda_3 + 1\lambda_4 + 4\lambda_5 + s_2^- = 1 \theta^* = 1 \\ & 2 \lambda_1 + 2\lambda_2 + 2\lambda_3 + 2\lambda_4 + 2\lambda_5 - s_1^+ = 2 \\ & \lambda_1, \lambda_2, \lambda_3, \lambda_4, \lambda_5, s_1^-, s_2^-, s_1^+ \geq 0 \end{aligned}$$

The optimal slacks are $s_1^{*-} = 2$, $s_2^{*-} = s_1^{*+} = 0$, with $\lambda_3^* = 1$ and all other $\lambda_j^* = 0$.

We now have

Definition 1 (DEA Efficiency): The performance of DMU_o is fully (100%) efficient if and only if both (i) $\theta^* = 1$ and (ii) all slacks $s_i^{*-} = s_r^{*+} = 0$.

Definition 2 (Weakly DEA Efficient): The performance of DMU_o is weakly efficient if and only if both (i) $\theta^* = 1$ and (ii) $s_i^{*-} \neq 0$ and/or $s_r^{*+} \neq 0$ for some i and r in some alternate optima.

Model (4) is usually called the second stage calculation of an envelopment model. In fact, the envelopment model can be written as:

$$\begin{aligned}
 & \min \theta - \varepsilon \left(\sum_{i=1}^m s_i^- + \sum_{r=1}^s s_r^+ \right) \\
 & \text{subject to} \\
 & \sum_{j=1}^n x_{ij} \lambda_j + s_i^- = \theta x_{io} \quad i = 1, 2, \dots, m; \\
 & \sum_{j=1}^n y_{rj} \lambda_j - s_r^+ = y_{ro} \quad r = 1, 2, \dots, s; \\
 & \lambda_j, s_i^-, s_r^+ \geq 0 \quad \forall i, j, r
 \end{aligned} \tag{5}$$

where the s_i^- and s_r^+ are slack variables used to convert the inequalities in (3) to equivalent equations. This is equivalent to solving (5) in two stages by first minimizing θ , then fixing $\theta = \theta^*$ as in (4), where the slacks are to be maximized without altering the previously determined value of $\theta = \theta^*$. Formally, this is equivalent to granting “preemptive priority” to the determination of θ^* in (3). In this manner, the fact that the non-Archimedean element ε is defined to be smaller than any positive real number is accommodated without having to specify the value of ε (Cooper, Seiford and Zhu, 2004).

The above models are called input-oriented DEA models, as possible input reductions are of interest while the outputs are kept at their current levels. Similarly, output-oriented models can be developed. These models focus on possible output increases while the inputs are kept at their current levels. The interested reader should refer to Cooper, Seiford and Zhu (2004).

The models in Table 1-2 are also known as CRS (constant returns to scale) models. If the constraint $\sum_{j=1}^n \lambda_j = 1$ is adjoined, they are referred to as variable returns to scale (VRS) models (Banker, Charnes, Cooper, 1984). This is due to the fact that $\sum_{j=1}^n \lambda_j = 1$ changes the shape of DEA frontier, and is related to the concept of returns to scale.

Table 1-2. CCR DEA Model

Input-oriented	
Envelopment model	Multiplier model
$\min \theta - \varepsilon \left(\sum_{i=1}^m s_i^- + \sum_{r=1}^s s_r^+ \right)$	$\max z = \sum_{r=1}^s \mu_r y_{ro}$
subject to	subject to
$\sum_{j=1}^n x_{ij} \lambda_j + s_i^- = \theta x_{io} \quad i = 1, 2, \dots, m;$	$\sum_{r=1}^s \mu_r y_{rj} - \sum_{i=1}^m v_i x_{ij} \leq 0$
$\sum_{j=1}^n y_{rj} \lambda_j - s_r^+ = y_{ro} \quad r = 1, 2, \dots, s;$	$\sum_{i=1}^m v_i x_{io} = 1$
$\lambda_j \geq 0 \quad j = 1, 2, \dots, n.$	$\mu_r, v_i \geq \varepsilon > 0$
Output-oriented	
Envelopment model	Multiplier model
$\max \phi + \varepsilon \left(\sum_{i=1}^m s_i^- + \sum_{r=1}^s s_r^+ \right)$	$\min q = \sum_{i=1}^m v_i x_{io}$
subject to	subject to
$\sum_{j=1}^n x_{ij} \lambda_j + s_i^- = x_{io} \quad i = 1, 2, \dots, m;$	$\sum_{i=1}^m v_i x_{ij} - \sum_{r=1}^s \mu_r y_{rj} \geq 0$
$\sum_{j=1}^n y_{rj} \lambda_j - s_r^+ = \phi y_{ro} \quad r = 1, 2, \dots, s;$	$\sum_{r=1}^s \mu_r y_{ro} = 1$
$\lambda_j \geq 0 \quad j = 1, 2, \dots, n.$	$\mu_r, v_i \geq \varepsilon > 0$

3. DATA AND STRUCTURE ISSUES

The current volume deals with data irregularities and structural complexities in applications of DEA.

Chapter 2 (by Cook and Zhu) develops a general framework for modeling and treating qualitative data in DEA and provides a unified structure for embedding rank order data into the DEA framework. It is shown that the existing approaches for dealing with qualitative data are equivalent.

Chapter 3 (by Chen and Zhu) discusses how to use the standard DEA models to deal with imprecise data in DEA (IDEA), concentrating on interval and ordinal data. There are two approaches in dealing with such imprecise inputs and outputs. One uses scale transformations and variable alternations to convert the non-linear IDEA model into a linear program, while the other identifies a set of exact data from the imprecise inputs and outputs and then uses the standard linear DEA model. The chapter focuses on the latter IDEA approach that uses the standard DEA model. It is shown that different results are obtained depending on whether the imprecise data are introduced directly into the multiplier or envelopment DEA model; the presence of imprecise data invalidates the linear duality between the multiplier and envelopment DEA models. The approaches are illustrated with both numerical and real world data sets.

Chapter 4 (by Pastor and Ruiz) presents an overview of the different existing approaches dealing with the treatment of negative data in DEA. Both the classical approaches and the most recent contributions to this problem are presented. The focus is mainly on issues such as translation invariance and units invariance of the variables, classification invariance of the units, as well as efficiency measurement and target setting.

Chapter 5 (by Ruggiero) discusses existing approaches to measuring performance when non-discretionary inputs affect the transformation of discretionary inputs into outputs. The suitability of the approaches depends on underlying assumptions about the relationship between non-discretionary inputs and outputs. One model treats non-discretionary inputs like discretionary inputs but uses a non-radial approach to project inefficient decision making units (DMUs) to the frontier holding non-discretionary inputs fixed. Other approaches use multiple stage models with regression to control for the effect that non-discretionary inputs have on production.

Chapter 6 (by Hua and Bian) discusses the existing methods of treating undesirable factors in DEA. Under strongly disposable technology and weakly disposable technology, there are several approaches for treating undesirable outputs in the DEA literature. One such approach is the hyperbolic output efficiency measure that increases desirable outputs and decreases undesirable outputs simultaneously. Based on the classification invariance property, a linear monotone decreasing transformation is used to treat the undesirable outputs. A directional distance function is used to estimate the efficiency scores based on weak disposability of undesirable outputs. This chapter also presents an extended DEA model in which

undesirable outputs and non-discretionary inputs are considered simultaneously.

Chapter 7 (by Piot-Lepetit and Le Moing) highlights the usefulness of the directional distance function in measuring the impact of the EU Nitrate directive, which prevents the free disposal of organic manure and nitrogen surplus. Efficiency indices for the production and environmental performance of farms at an individual level are proposed, together with an evaluation of the impact caused by the said EU regulation. This chapter extends the previous approach to good and bad outputs within the framework of the directional distance function, by introducing a by-product (organic manure), which becomes a pollutant once a certain level of disposability is exceeded.

Chapter 8 (by Adler and Golany) presents the combined use of principal component analysis (PCA) and DEA with the stated aim of reducing the curse of dimensionality that occurs in DEA when there is an excessive number of inputs and outputs in relation to the number of decision-making units. Various PCA-DEA formulations are developed in the chapter utilizing the results of principal component analyses to develop objective assurance region type constraints on the DEA weights. The first set of models applies PCA to grouped data representing similar themes, such as quality or environmental measures. The second set of models, if needed, applies PCA to all inputs and separately to all outputs, thus further strengthening the discrimination power of DEA. A case study of municipal solid waste managements in the Oulu district of Finland, which has been frequently analyzed in the literature, will illustrate the different models and the power of the PCA-DEA formulation.

Chapter 9 (by Dulá) deals with the extension of data envelopment analysis to the general problem of mining oriented outliers. DEA is firmly anchored in efficiency and productivity paradigms. This research claims new application domains for DEA by releasing it from these moorings. The same reasons why efficient entities are of interest in DEA apply to the geometric equivalent in general point sets since they are based on the data's magnitude limits relative to the other data points. A framework for non-parametric frontier analysis is derived from a new set of first principles.

Chapter 10 (by Adler, Raveh and Yazhemsky) presents the results of DEA in a two-dimensional plot. Presenting DEA graphically, due to its multiple variable nature, has proven difficult and some have argued that this has left decision-makers at a loss in interpreting the results. Co-Plot, a

variant of multi-dimensional scaling, locates each decision-making unit in a two-dimensional space in which the location of each observation is determined by all variables simultaneously. The graphical display technique exhibits observations as points and variables (ratios) as arrows, relative to the same center-of-gravity. Observations are mapped such that similar decision-making units are closely located on the plot, signifying that they belong to a group possessing comparable characteristics and behavior.

Chapter 11 (by Cook, Liang, Yang and Zhu) presents several DEA-based approaches for characterizing and measuring supply chain efficiency. The models are illustrated in a seller-buyer supply chain context, when the relationship between the seller and buyer is treated leader-follower and cooperative, respectively. In the leader-follower structure, the leader is first evaluated, and then the follower is evaluated using information related to the leader's efficiency. In the cooperative structure, the joint efficiency which is modeled as the average of the seller's and buyer's efficiency scores is maximized, and both supply chain members are evaluated simultaneously.

Chapter 12 (by Färe, Grosskopf and Whittaker) describes network DEA models, where a network consists of sub-technologies. A DEA model typically describes a technology to a level of abstraction necessary for the analyst's purpose, but leaves out a description of the sub-technologies that make up the internal functions of the technology. These sub-technologies are usually treated as a "black box", i.e., there is no information about what happens inside them. The specification of the sub-technologies enables the explicit examination of input allocation and intermediate products that together form the production process. The combination of sub-technologies into networks provides a method of analyzing problems that the traditional DEA models cannot address.

Chapter 13 (Morita and Zhu) presents a context-dependent DEA methodology, which refers to a DEA approach where a set of DMUs is evaluated against a particular evaluation context. Each evaluation context represents an efficient frontier composed of DMUs in a specific performance level. The context-dependent DEA measures the attractiveness and the progress when DMUs exhibiting poorer and better performance are chosen as the evaluation context, respectively. This chapter also presents a slack-based context-dependent DEA approach. In DEA, nonzero input and output slacks are very likely to be present, after the radial efficiency score improvement. Slack-based context-dependent DEA allows us to fully evaluate the inefficiency in a DMU's performance.

Chapter 14 (by Cook and Zhu) presents DEA models to accommodate flexible measures. In standard DEA, it is assumed that the input versus output status of each of the chosen analysis measures is known. In some situations, however, certain measures can play either input or output roles. Consider using the number of nurse trainees on staff in a study of hospital efficiency. Such a factor clearly constitutes an output measure for a hospital, but at the same time is an important component of the hospital's total staff complement, hence is an input. Both an individual DMU model and an aggregate model are suggested as methodologies for deriving the most appropriate designations for flexible measures.

Chapter 15 (by Lozano and Villa) presents DEA models under situations where one or more inputs and/or outputs are integer quantities. Commonly, in these situations, the non-integer targets are rounded off. However, rounding off may easily lead to an infeasible target (i.e. out of the Production Possibility Set) or to a dominated operation point. In this chapter, a general framework to handle integer inputs and outputs is presented and a number of integer DEA models are reviewed.

Chapter 16 (by Sarkis) looks at some data requirements and characteristics that may ease the execution of the DEA models and the interpretation of DEA results.

REFERENCES

1. Banker, R., A. Charnes and W.W. Cooper, 1984, Some models for estimating technical and scale inefficiencies in data envelopment analysis, *Management Science* 30, 1078-1092.
2. Charnes A., W.W. Cooper and E. Rhodes, 1978, Measuring the efficiency of decision making units, *European Journal of Operational Research*, 2(6), 428-444.
3. Cooper, W.W., L.M. Seiford and J. Zhu, 2004, *Handbook of Data Envelopment Analysis*, Kluwer Academic Publishers, Boston.
4. Zhu, J. 2003, *Quantitative Models for Performance Evaluation and Benchmarking: Data Envelopment Analysis with Spreadsheets and DEA Excel Solver*, Kluwer Academic Publishers, Boston.

Chapter 2

RANK ORDER DATA IN DEA

Wade D. Cook¹ and Joe Zhu²

¹*Schulich School of Business, York University, Toronto, Ontario, Canada, M3J 1P3,
wcook@shulich.yorku.ca*

²*Department of Management, Worcester Polytechnic Institute, Worcester, MA 01609,
jzhu@wpi.edu*

Abstract: In data envelopment analysis (DEA), performance evaluation is generally assumed to be based upon a set of quantitative data. In many real world settings, however, it is essential to take into account the presence of qualitative factors when evaluating the performance of decision making units (DMUs). Very often rankings are provided from best to worst relative to particular attributes. Such rank positions might better be presented in an ordinal, rather than numerical sense. The chapter develops a general framework for modeling and treating qualitative data in DEA, and provides a unified structure for embedding rank order data into the DEA framework. We show that the approach developed earlier in Cook et al (1993, 1996) is equivalent to the IDEA methodology given in Chapter 3. It is shown that, like IDEA, the approach given here for dealing with qualitative data lends itself to treatment by conventional DEA methodology.

Key words: Data envelopment analysis (DEA), efficiency, qualitative, rank order data

1. INTRODUCTION

In the data envelopment analysis (DEA) model of Charnes, Cooper and Rhodes (1978), each member of a set of n decision making units (DMUs) is to be evaluated relative to its peers. This evaluation is generally assumed to be based on a set of *quantitative* output and input factors. In many real world settings, however, it is essential to take into account the presence of

qualitative factors when rendering a decision on the performance of a DMU. Very often it is the case that for a factor such as management competence, one can, at most, provide a *ranking* of the DMUs from best to worst relative to this attribute. The capability of providing a more precise, quantitative measure reflecting such a factor is often not feasible. In some situations such factors can be legitimately “quantified,” but very often such quantification may be superficially forced as a modeling convenience.

In situations such as that described, the “data” for certain influence factors (inputs and outputs) might better be represented as rank positions in an ordinal, rather than numerical sense. Refer again to the management competence example. In certain circumstances, the information available may only permit one to put each DMU into one of L categories or groups (e.g. ‘high’, ‘medium’ and ‘low’ competence). In other cases, one may be able to provide a complete rank ordering of the DMUs on such a factor.

Cook, Kress and Seiford (1993), (1996) first presented a modified DEA structure, incorporating rank order data. The 1996 article applied this structure to the problem of prioritizing a set of research and development projects, where there were both inputs and outputs defined on a Likert scale. An alternative to the Cook et al approach was provided in Cooper, Park and Yu (1999) in the form of the imprecise DEA (IDEA) model. While various forms of imprecise data were examined, one major component of that research focused on ordinal (rank order) data. See Chapter 3 for a treatment of the specifics of IDEA. These two approaches to the treatment of ordinal data in DEA are further discussed and compared in Cook and Zhu(2006).

In the current chapter, we present a unified structure for embedding rank order or Likert scale data into the DEA framework. This development is very much related to the presentation in Cook and Zhu (2006). To provide a practical setting for the methodology to be developed herein, Section 2 briefly discusses the R&D project selection problem as presented in more detail in Cook et al (1996). Section 3 presents a continuous projection model, based on the conventional radial model of Charnes et al (1978). In Section 4 this approach is compared to the IDEA methodology of Cooper et al (1999). We demonstrate that IDEA for Likert Scale data is in fact equivalent to the earlier approach of Cook, Kress and Seiford (1996). Section 5 develops a discrete projection methodology that guarantees projection to points on the Likert Scale. Conclusions and further directions are addressed in Section 6.

2. ORDINAL DATA IN R&D PROJECT SELECTION

Consider the problem of selecting R&D projects in a major public utility corporation with a large research and development branch. Research activities are housed within several different divisions, for example, thermal, nuclear, electrical, and so on. In a budget constrained environment in which such an organization finds itself, it becomes necessary to make choices among a set of potential research initiatives or projects that are in competition for the limited resources. To evaluate the impact of funding (or not funding) any given research initiative, two major considerations generally must be made. First, the initiative must be viewed in terms of more than one factor or criterion. Second, some or all of the criteria that enter the evaluation may be qualitative in nature. Even when pure quantitative factors are involved, such as long term saving to the organization, it may be difficult to obtain even a crude estimate of the value of that factor. The most that one can do in many such situations is to classify the project (according to this factor) on some scale (high/medium/low or say a 5-point scale).

Let us assume that for each qualitative criterion each initiative is rated on a 5-point scale, where the particular point on the scale is chosen through a consensus on the part of executives within the organization. Table 2-1 presents an illustration of how the data might appear for 10 projects, 3 qualitative output criteria (benefits), and 3 qualitative input criteria (cost or resources). In the actual setting examined, a number of potential benefit and cost criteria were considered as discussed in Cook et al (1996).

We use the convention that for both outputs and inputs, a rating of 1 is “best,” and 5 “worst.” For outputs, this means that a DMU ranked at position 1 generates *more* output than is true of a DMU in position 2, and so on. For inputs, a DMU in position 1 consumes *less* input than one in position 2.

Table 2-1. Ratings by criteria

Project No.	Outputs			Inputs		
	1	2	3	4	5	6
1	2	4	1	5	2	1
2	1	1	4	3	5	2
3	1	1	1	1	2	1
4	3	3	3	4	3	2
5	4	3	5	5	1	4
6	2	5	1	1	2	2
7	1	4	1	5	4	3
8	1	5	3	3	3	3
9	5	2	4	4	2	5
10	5	4	4	5	5	5

Regardless of the manner in which such a scale rating is arrived at, the existing DEA model is capable only of treating the information as if it has cardinal meaning (e.g. something which receives a score of 4 is evaluated as being twice as important as something that scores 2). There are a number of problems with this approach. First and foremost, the projects' original data in the case of some criteria may take the form of an ordinal ranking of the projects. Specifically, the most that can be said about two projects i and j is that i is preferred to j . In other cases it may only be possible to classify projects as say 'high', 'medium' or 'low' in importance on certain criteria. When projects are rated on, say, a 5-point scale, it is generally understood that this scale merely provides a relative positioning of the projects. In a number of agencies investigated (for example, hydro electric and telecommunications companies), 5-point scales are common for evaluating alternatives in terms of qualitative data, and are often accompanied by interpretations such as

- 1 = Extremely important
- 2 = Very important
- 3 = Important
- 4 = Low in importance
- 5 = Not important,

that are easily understood by management. While it is true that market researchers often treat such scales in a numerical (i.e. cardinal) sense, it is not practical that in rating a project, the classification 'extremely important' should be interpreted literally as meaning that this project rates three times better than one which is only classified as 'important.' The key message here is that many, if not all criteria used to evaluate R&D projects are qualitative in nature, and should be treated as such. The model presented in the following sections extends the DEA idea to an ordinal setting, hence accommodating this very practical consideration.

3. MODELING LIKERT SCALE DATA: CONTINUOUS PROJECTION

The above problem typifies situations in which pure ordinal data or a mix of ordinal and numerical data, are involved in the performance measurement exercise. To cast this problem in a general format, consider the situation in which a set of N decision making units (DMUs), $k=1,\dots,N$ are to be evaluated in terms of R_1 numerical outputs, R_2 ordinal outputs, I_1 numerical inputs, and I_2 ordinal inputs. Let $Y_k^1 = (y_{rk}^1)$, $Y_k^2 = (y_{rk}^2)$ denote the R_1 -

dimensional and R_2 -dimensional vectors of outputs, respectively. Similarly, let $X_k^1 = (x_{ik}^1)$ and $X_k^2 = (x_{ik}^2)$ be the I_1 and I_2 -dimensional vectors of inputs, respectively.

In the situation where all factors are quantitative, the conventional radial projection model for measuring the efficiency of a DMU is expressed by the ratio of weighted outputs to weighted inputs. Adopting the general variable returns to scale (VRS) model of Banker, Charnes and Cooper (1984), the efficiency of DMU "0" follows from the solution of the optimization model:

$$\begin{aligned} e_0 &= \max \left(\mu_0 + \sum_{r \in R_1} \mu_r^1 y_{ro}^1 + \sum_{r \in R_2} \mu_r^2 y_{ro}^2 \right) / \left(\sum_{i \in I_1} v_i^1 x_{io}^1 + \sum_{i \in I_2} v_i^2 x_{io}^2 \right) \\ \text{s.t.} \\ &(\mu_0 + \sum_{r \in R_1} \mu_r^1 y_{rk}^1 + \sum_{r \in R_2} \mu_r^2 y_{rk}^2) / \left(\sum_{i \in I_1} v_i^1 x_{ik}^1 + \sum_{i \in I_2} v_i^2 x_{ik}^2 \right) \leq 1, \text{ all } k \quad (3.1) \\ &\mu_r^1, \mu_r^2, v_i^1, v_i^2 \geq \varepsilon, \text{ all } r, i \end{aligned}$$

Problem (3.1) is convertible to the linear programming format:

$$\begin{aligned} e_0 &= \max \mu_0 + \sum_{r \in R_1} \mu_r^1 y_{ro}^1 + \sum_{r \in R_2} \mu_r^2 y_{ro}^2 \\ \text{s.t.} \quad &\sum_{i \in I_1} v_i^1 x_{io}^1 + \sum_{i \in I_2} v_i^2 x_{io}^2 = 1 \quad (3.2) \\ &\mu_0 + \sum_{r \in R_1} \mu_r^1 y_{rk}^1 + \sum_{r \in R_2} \mu_r^2 y_{rk}^2 - \sum_{i \in I_1} v_i^1 x_{ik}^1 - \sum_{i \in I_2} v_i^2 x_{ik}^2 \leq 0, \text{ all } k \\ &\mu_r^1, \mu_r^2, v_i^1, v_i^2 \geq \varepsilon, \text{ all } r, i, \end{aligned}$$

whose dual is given by

$$\begin{aligned} \min \quad &\theta - \varepsilon \sum_{r \in R_1 \cup R_2} s_r^+ - \varepsilon \sum_{i \in I_1 \cup I_2} s_i^- \\ \text{s.t.} \quad &\sum_{k=1}^N \lambda_k y_{rk}^1 - s_r^+ = y_{ro}^1, \quad r \in R_1 \\ &\sum_{k=1}^N \lambda_k y_{rk}^2 - s_r^+ = y_{ro}^2, \quad r \in R_2 \\ &\theta x_{io}^1 - \sum_{k=1}^N \lambda_k x_{ik}^1 - s_i^- = 0, \quad i \in I_1 \\ &\theta x_{io}^2 - \sum_{k=1}^N \lambda_k x_{ik}^2 - s_i^- = 0, \quad i \in I_2 \\ &\sum_{k=1}^N \lambda_k = 1 \\ &\lambda_k, s_r^+, s_i^- \geq 0, \text{ all } k, r, i, \theta \text{ unrestricted} \end{aligned} \quad (3.2')$$

To place the problem in a general framework, assume that for each ordinal factor ($r \in R_2$, $i \in I_2$), a DMU k can be assigned to one of L rank positions, where $L \leq N$. As discussed earlier, $L=5$ is an example of an appropriate number of rank positions in many practical situations. We point out that in certain application settings, different ordinal factors may have different L -values associated with them. For exposition purposes, we assume a common L -value throughout. We demonstrate later that this represents no loss of generality.

One can view the allocation of a DMU to a rank position ℓ on an output r , for example, as having assigned that DMU an output *value* or *worth* $y_r^2(\ell)$. The implementation of the DEA model (3.1) (and (3.2)) thus involves determining two things:

- (1) multiplier values μ_r^2 , ν_i^2 for outputs $r \in R_2$ and inputs $i \in I_2$;
- (2) rank position values $y_r^2(\ell)$, $r \in R_2$, and $x_i^2(\ell)$, $i \in I_2$, all ℓ .

In this section we show that the problem can be reduced to the standard VRS model by considering (1) and (2) simultaneously.

To facilitate development herein, define the L -dimensional unit vectors $\gamma_{rk} = (\gamma_{rk}(\ell))$, and $\delta_{ik} = (\delta_{ik}(\ell))$ where

$$\gamma_{rk}(\ell) = \begin{cases} 1 & \text{if DMU } k \text{ is ranked in } \ell \text{ th position on output } r \\ 0, & \text{otherwise} \end{cases}$$

$$\delta_{ik}(\ell) = \begin{cases} 1 & \text{if DMU } k \text{ is ranked in } \ell \text{ th position on input } i \\ 0, & \text{otherwise} \end{cases}$$

For example, if a 5-point scale is used, and if DMU #1 is ranked in $\ell = 3^{\text{rd}}$ place on ordinal output $r=5$, then $\gamma_{51}(3)=1$, $\gamma_{51}(\ell)=0$, for all other rank positions ℓ . Thus, y_{51}^2 is assigned the value $y_5^2(3)$, the *worth* to be credited to the 3^{rd} rank position on output factor 5. It is noted that y_{rk}^2 can be represented in the form

$$y_{rk}^2 = y_r^2(\ell_{rk}) = \sum_{\ell=1}^L y_r^2(\ell) \gamma_{rk}(\ell),$$

where ℓ_{rk} is the rank position occupied by DMU k on output r . Hence, model (3.2) can be rewritten in the more representative format.

$$\begin{aligned} e_o = \max \quad & \mu_o + \sum_{r \in R_1} \mu_r^1 y_{ro}^1 + \sum_{r \in R_2} \sum_{\ell=1}^L \mu_r^2 y_r^2(\ell) \gamma_{ro}(\ell) \\ \text{s.t.} \quad & \sum_{i \in I_1} \nu_i^1 x_{io}^1 + \sum_{i \in I_2} \sum_{\ell=1}^L \nu_i^2 x_i^2(\ell) \delta_{io}(\ell) = 1 \\ & \mu_o + \sum_{r \in R_1} \mu_r^1 y_{rk}^1 + \sum_{r \in R_2} \sum_{\ell=1}^L \mu_r^2 y_r^2(\ell) \gamma_{rk}(\ell) - \sum_{i \in I_1} \nu_i^1 x_{ik}^1 - \sum_{i \in I_2} \sum_{\ell=1}^L \nu_i^2 x_i^2(\ell) \delta_{ik}(\ell) \\ & \leq 0, \text{ all } k \end{aligned} \tag{3.3}$$

$$\{Y_r^2 = (y_r^2(\ell)), X_i^2 = (x_i^2(\ell))\} \in \Psi$$

$$\mu_r^1, \nu_i^1 \geq \varepsilon$$

In (3.3) we use the notation Ψ to denote the set of *permissible worth vectors*. We discuss this set below.

It must be noted that the same infinitesimal ε is applied here for the various input and output multipliers, which may, in fact, be measured on scales that are very different from another. If two inputs are, for example, x_{i1k}^1 representing ‘labor hours’, and x_{i2k}^1 representing ‘available computer technology’, the scales would clearly be incompatible. Hence, the likely sizes of the corresponding multipliers ν_{i1}^1, ν_{i2}^1 may be correspondingly different. Thrall (1996) has suggested a mechanism for correcting for such scale incompatibility, by applying a *penalty vector* G to augment ε , thereby creating differential lower bounds on the various ν_i, μ_r . Proper choice of G can effectively bring all factors to some form of common scale or unit. For simplicity of presentation we will assume the cardinal scales for all $r \in R_1, i \in I_1$ are similar in dimension, and that G is the unit vector. The more general case would proceed in an analogous fashion.

Permissible Worth Vectors

The values or worths $\{y_r^2(\ell)\}, \{x_i^2(\ell)\}$, attached to the ordinal rank positions for outputs r and inputs i , respectively, must satisfy the minimal requirement that it is *more* important to be ranked in ℓ^{th} position than in the $(\ell+1)^{\text{st}}$ position on any such ordinal factor. Specifically, $y_r^2(\ell) > y_r^2(\ell+1)$ and $x_i^2(\ell) < x_i^2(\ell+1)$. That is, for outputs, one places a higher weight on being ranked in ℓ^{th} place than in $(\ell+1)^{\text{st}}$ place. For inputs, the opposite is true. A set of linear conditions that produce this realization is defined by the set Ψ , where

$$\Psi = \{(Y_r^2, X_i^2) | y_r^2(\ell) - y_r^2(\ell+1) \geq \varepsilon, \ell=1, \dots, L-1, y_r^2(L) \geq \varepsilon, \\ x_i^2(\ell+1) - x_i^2(\ell) \geq \varepsilon, \ell=1, \dots, L-1, x_i^2(1) \geq \varepsilon\}.$$

Arguably, ε could be made dependent upon ℓ (i.e. replace ε by ε_ℓ). It can be shown, however, that all results discussed below would still follow. For convenience, we, therefore, assume a common value for ε . We now demonstrate that the nonlinear problem (3.3) can be written as a linear programming problem.

Theorem 3.1

Problem (3.3), in the presence of the permissible worth space Ψ , can be expressed as a linear programming problem.

Proof: In (3.3), make the change of variables $w_{r\ell}^1 = \mu_r^2 y_r^2(\ell)$, $w_{i\ell}^2 = \nu_i^2 x_i^2(\ell)$.

It is noted that in Ψ , the expressions $y_r^2(\ell) - y_r^2(\ell+1) \geq \varepsilon$, $y_r^2(L) \geq \varepsilon$ can be replaced by

$\mu_r^2 y_r^2(\ell) - \mu_r^2 y_r^2(\ell+1) \geq \mu_r^2 \varepsilon$, $\mu_r^2 y_r^2(L) \geq \mu_r^2 \varepsilon$, which becomes $w_{r\ell}^1 - w_{r\ell+1}^1 \geq \mu_r^2 \varepsilon$, $w_{rL}^1 \geq \mu_r^2 \varepsilon$.

A similar conversion holds for the $x_i^2(\ell)$. Problem (3.3) now becomes

$$\begin{aligned}
 e_0 = \max \quad & \mu_o + \sum_{r \in R_1} \mu_r^1 y_{ro}^1 + \sum_{r \in R_2} \sum_{\ell=1}^L w_{r\ell}^1 \gamma_{ro}(\ell) \\
 \text{s.t.} \quad & \sum_{i \in I_1} \nu_i^1 x_{io}^1 + \sum_{i \in I_2} \sum_{\ell=1}^L w_{i\ell}^2 \delta_{io}(\ell) = 1 \\
 & \mu_o + \sum_{r \in R_1} \mu_r^1 y_{rk}^1 + \sum_{r \in R_2} \sum_{\ell=1}^L w_{r\ell}^1 \gamma_{rk}(\ell) \\
 & - \sum_{i \in I_1} \nu_i^1 x_{ik}^1 - \sum_{i \in I_2} \sum_{\ell=1}^L w_{i\ell}^2 \delta_{ik}(\ell) \leq 0, \text{ all } k \\
 & w_{r\ell}^1 - w_{r\ell+1}^1 \geq \mu_r^2 \varepsilon, \ell=1, \dots, L-1, \text{ all } r \in R_2 \\
 & w_{rL}^1 \geq \mu_r^2 \varepsilon, \text{ all } r \in R_2 \\
 & w_{i\ell+1}^2 - w_{i\ell}^2 \geq \nu_i^2 \varepsilon, \ell=1, \dots, L-1, \text{ all } i \in I_2 \\
 & w_{iL}^2 \geq \nu_i^2 \varepsilon, \text{ all } i \in I_2 \\
 & \mu_r^1, \nu_i^1 \geq \varepsilon, \text{ all } r \in R_1, i \in I_1 \\
 & \mu_r^2, \nu_i^2 \geq \varepsilon, \text{ all } r \in R_2, i \in I_2
 \end{aligned} \tag{3.4}$$

Problem (3.4) is clearly in linear programming problem format.

QED

We state without proof the following theorem.

Theorem 3.2

At the optimal solution to (3.4), $\mu_r^2 = \nu_i^2 = \varepsilon$ for all $r \in R_2, i \in I_2$.

Problem (3.4) can then be expressed in the form:

$$\begin{aligned}
 e_0 = \max \quad & \mu_o + \sum_{r \in R_1} \mu_r^1 y_{ro}^1 + \sum_{r \in R_2} \sum_{\ell=1}^L w_{r\ell}^1 \gamma_{ro}(\ell) \\
 \text{s.t.} \quad & \sum_{i \in I_1} \nu_i^1 x_{io}^1 + \sum_{i \in I_2} \sum_{\ell=1}^L w_{i\ell}^2 \delta_{io}(\ell) = 1
 \end{aligned}$$

$$\begin{aligned}
& \mu_o + \sum_{r \in R_1} \mu_r^1 y_{rk}^1 + \sum_{r \in R_2} \sum_{\ell=1}^L w_{r\ell}^1 \gamma_{rk}(\ell) - \sum_{i \in I_1} v_i^1 x_{ik}^1 - \sum_{i \in I_2} \sum_{\ell=1}^L w_{i\ell}^2 \delta_{ik} \\
& (\ell) \leq 0, \text{ all } k \\
& -w_{r\ell}^1 + w_{r\ell+1}^1 \leq -\varepsilon^2, \ell=1, \dots, L-1, \text{ all } r \in R_2 \\
& -w_{rL}^1 \leq -\varepsilon^2, \text{ all } r \in R_2 \\
& -w_{i\ell+1}^2 + w_{i\ell}^2 \leq -\varepsilon^2, \ell=1, \dots, L-1, \text{ all } i \in I_2 \\
& -w_{i1}^2 \leq -\varepsilon^2, \text{ all } i \in I_2 \\
& \mu_r^1, v_i^1 \geq \varepsilon, r \in R_1, i \in I_1
\end{aligned} \tag{3.5}$$

It can be shown that (3.5) is equivalent to the *standard* VRS model. First we form the dual of (3.5).

$$\min \theta - \varepsilon \sum_{r \in R_1} s_r^+ - \varepsilon \sum_{i \in I_1} s_i^- - \varepsilon^2 \sum_{r \in R_2} \sum_{\ell=1}^L \alpha_{r\ell}^1 - \varepsilon^2 \sum_{i \in I_2} \sum_{\ell=1}^L \alpha_{i\ell}^2$$

$$\text{s.t.} \quad \sum_{k=1}^N \lambda_k y_{rk}^1 - s_r^+ = y_{ro}^1, r \in R_1$$

$$\theta x_{io}^1 - \sum_{k=1}^N \lambda_k x_{ik}^1 - s_i^- = 0, i \in I_1$$

(3.5')

$$\left. \begin{aligned}
& \sum_{k=1}^N \lambda_k \gamma_{rk}(1) - \alpha_{r1}^1 = \gamma_{ro}(1) \\
& \sum_{k=1}^N \lambda_k \gamma_{rk}(2) + \alpha_{r1}^1 - \alpha_{r2}^1 = \gamma_{ro}(2) \\
& \vdots \\
& \sum_{k=1}^N \lambda_k \gamma_{rk}(L) + \alpha_{rL-1}^1 - \alpha_{rL}^1 = \gamma_{ro}(L)
\end{aligned} \right\} r \in R_2$$

$$\left. \begin{aligned}
& \delta_{io}(L)\theta - \sum_{k=1}^N \lambda_k \delta_{ik}(L) - \alpha_{iL}^2 = 0 \\
& \delta_{io}(L-1)\theta - \sum_{k=1}^N \lambda_k \delta_{ik}(L-1) + \alpha_{iL}^2 - \alpha_{iL-1}^2 = 0 \\
& \vdots \\
& \delta_{io}(1)\theta - \sum_{k=1}^N \lambda_k \delta_{ik}(1) + \alpha_{i2}^2 - \alpha_{i1}^2 = 0
\end{aligned} \right\} i \in I_2$$

$$\sum_{k=1}^N \lambda_k = 1$$

$$\lambda_k, s_r^+, s_i^-, \alpha_{r\ell}^1, \alpha_{i\ell}^2 \geq 0, \theta \text{ unrestricted.}$$

Here, we use $\{\lambda_k\}$ as the standard dual variables associated with the N ratio constraints, and the variables $\{\alpha_{i\ell}^2, \alpha_{r\ell}^1\}$ are the dual variables associated

with the rank order constraints defined by Ψ . The slack variables s_r^+ , s_i^- correspond to the lower bound restrictions on μ_r^1 , v_i^1 .

Now, perform simple row operations on (3.5') by replacing the ℓ^{th} constraint by the sum of the first ℓ constraints. That is, the second constraint (for those $r \in R_2$ and $i \in I_2$) is replaced by the sum of the first two constraints, constraint 3 by the sum of the first three, and so on. Letting

$$\bar{\gamma}_{rk}(\ell) = \sum_{n=1}^{\ell} \gamma_{rk}(n) = \gamma_{rk}(1) + \gamma_{rk}(2) + \dots + \gamma_{rk}(\ell),$$

and

$$\bar{\delta}_{ik}(\ell) = \sum_{n=\ell}^L \delta_{ik}(n) = \delta_{ik}(L) + \delta_{ik}(L-1) + \dots + \delta_{ik}(\ell),$$

problem (3.5') can be rewritten as:

$$\begin{aligned} \min \quad & \theta - \varepsilon \sum_{r \in R_1} s_r^+ - \varepsilon \sum_{i \in I_1} s_i^- - \varepsilon^2 \sum_{r \in R_2} \sum_{\ell=1}^L \alpha_{r\ell}^1 - \varepsilon^2 \sum_{i \in I_2} \sum_{\ell=1}^L \alpha_{i\ell}^2 \\ \text{s.t.} \quad & \sum_{k=1}^N \lambda_k y_{rk}^1 - s_r^+ = y_{ro}^1, r \in R_1 \\ & \theta x_{io}^1 - \sum_{k=1}^N \lambda_k x_{ik}^1 - s_i^- = 0, i \in I_1 \\ & \sum_{k=1}^N \lambda_k \bar{\gamma}_{rk}(\ell) - \alpha_{r\ell}^1 = \bar{\gamma}_{ro}(\ell), r \in R_2, \ell=1, \dots, L \\ & \theta \bar{\delta}_{io}(\ell) - \sum_{k=1}^N \lambda_k \bar{\delta}_{ik}(\ell) - \alpha_{i\ell}^2 = 0, i \in I_2, \ell=1, \dots, L \\ & \sum_{k=1}^N \lambda_k = 1 \\ & \lambda_k, s_r^+, s_i^-, \alpha_{r\ell}^1, \alpha_{i\ell}^2 \geq 0, \text{ all } i, r, \ell, k, \theta \text{ unrestricted in sign} \end{aligned} \quad (3.6')$$

The dual of (3.6') has the format:

$$\begin{aligned} e_o = \max \quad & \mu_o + \sum_{r \in R_1} \mu_r^1 y_{ro}^1 + \sum_{r \in R_2} \sum_{\ell=1}^L \bar{w}_{r\ell}^1 \bar{\gamma}_{ro}(\ell) \\ \text{s.t.} \quad & \sum_{i \in I_1} v_i^1 x_{io}^1 + \sum_{i \in I_2} \sum_{\ell=1}^L \bar{w}_{i\ell}^2 \bar{\delta}_{io}(\ell) = 1 \\ & \mu_o + \sum_{r \in R_1} \mu_r^1 y_{rk}^1 + \sum_{r \in R_2} \sum_{\ell=1}^L \bar{w}_{r\ell}^1 \bar{\gamma}_{rk}(\ell) - \sum_{i \in I_1} v_i^1 x_{ik}^1 - \sum_{i \in I_2} \sum_{\ell=1}^L \bar{w}_{i\ell}^2 \bar{\delta}_{ik}(\ell) \leq 0, \text{ all } k \\ & \mu_r^1, v_i^1 \geq \varepsilon, \bar{w}_{r\ell}^1, \bar{w}_{i\ell}^2 \geq \varepsilon^2, \end{aligned} \quad (3.6)$$

which is a form of the VRS model. The slight difference between (3.6) and the conventional VRS model of Banker et al. (1984), is the presence of a different ε (i.e., ε^2) relating to the multipliers $w_{r\ell}^1$, $w_{i\ell}^2$, than is true for the multipliers μ_r^1 , v_i^1 . It is observed that in (3.6') the common L-value can easily be replaced by criteria specific values (e.g. L_r for output criterion r).

The model structure remains the same, as does that of model (3.6). Of course, since the intention is to have an infinitesimal lower bound on multipliers (i.e., $\varepsilon \approx 0$), one can, from the start, restrict

$$\mu_r^1, \nu_i^1 \geq \varepsilon^2 \text{ and } \mu_r^2, \nu_i^2 \geq \varepsilon.$$

This leads to a form of (3.6) where all multipliers have the same infinitesimal lower bounds, making (3.6) precisely a VRS model in the spirit of Banker et al. (1984).

Criteria Importance

The presence of ordinal data factors results in the need to *impute* values $y_r^2(\ell)$, $x_i^2(\ell)$ to outputs and inputs, respectively, for DMUs that are ranked at positions on an L-point Likert or ordinal scale. Specifically, all DMUs ranked at that position will be credited with the same “amount” $y_r^2(\ell)$ of output r ($r \in R_2$) and $x_i^2(\ell)$ of input i ($i \in I_2$).

A consequence of the change of variables undertaken above, to bring about linearization of the otherwise nonlinear terms, e.g., $w_{r\ell}^1 = \mu_r^2 y_r^2(\ell)$, is that at the optimum, all $\mu_r^2 = \varepsilon^2$, $\nu_i^2 = \varepsilon^2$. Thus, all of the ordinal criteria are relegated to the status of being of *equal importance*. Arguably, in many situations, one may wish to view the relative importance of these ordinal criteria (as captured by the μ_r^2 , ν_i^2) in the same spirit as we have viewed the data values $\{y_{rk}^2\}$. That is, there may be sufficient information to be able to *rank* these criteria. Specifically, suppose that the R_2 output criteria can be grouped into L_1 categories and the I_2 input criteria into L_2 categories. Now, replace the variables μ_r^2 by $\mu^2(m)$, and ν_i^2 by $\nu^2(n)$, and restrict:

$$\begin{aligned} \mu^2(m) - \mu^2(m+1) &\geq \varepsilon, m=1, \dots, L_1-1 \\ \mu^2(L_1) &\geq \varepsilon \\ \text{and} \\ \nu^2(n) - \nu^2(n+1) &\geq \varepsilon, n=1, \dots, L_2-1 \\ \nu^2(L_2) &\geq \varepsilon. \end{aligned}$$

Letting m_r denote the rank position occupied by output $r \in R_2$, and n_i the rank position occupied by input $i \in I_2$, we perform the change of variables

$$\begin{aligned} w_{r\ell}^1 &= \mu^2(m_r) y_r^2(\ell) \\ w_{i\ell}^2 &= \nu^2(n_i) x_i^2(\ell) \end{aligned}$$

The corresponding version of model (3.4) would see the lower bound restrictions μ_r^2 , $\nu_i^2 \geq \varepsilon$ replaced by the above constraints on $\mu^2(m)$ and $\nu^2(n)$.

(n). Again, arguing that at the optimum in (3.4), these variables will be forced to their lowest levels, the resulting values of the $\mu^2(m)$, $\nu^2(n)$ will be $\mu^2(m) = (L_1+1-m)\varepsilon$, $\nu^2(n) = (L_2+1-n)\varepsilon$.

This implies that the lower bound restrictions on $w_{r\ell}^1$, $w_{i\ell}^2$ become

$$w_{r\ell}^1 \geq (L_1+1-m_r)\varepsilon^2, w_{i\ell}^2 \geq (L_2+1-n_i)\varepsilon^2.$$

Example

When model (3.6') is applied to the data of Table 2-1, the efficiency scores obtained are as shown in Table 2-2.

Table 2-2. Efficiency Scores (Non-ranked Criteria)

Project	1	2	3	4	5	6	7	8	9	10
Score	0.76	0.73	1.00	0.67	1.00	0.82	0.67	0.67	0.55	0.37

Here, projects 3 and 5 turn out to be 'efficient', while all other projects are rated well below 100%. In this particular analysis, ε was chosen as 0.03. In another run (not shown here) where $\varepsilon = 0.01$ was used, projects 3, 5 and 6 received ratings of 1.00, while all others obtained somewhat higher scores than those shown in Table 2-2. When a very small value of ε ($\varepsilon=0.001$) was used, all except one of the projects was rated as efficient. Clearly this example demonstrates the same degree of dependence on the choice of ε as is true in the standard DEA model. See Ali and Seiford (1993).

From the data in Table 2-1 it might appear that only project 3 should be efficient since 3 dominates project 5 in all factors except for the second input where project 3 rates second while project 5 rates first. As is characteristic of the standard ratio DEA model, a single factor can produce such an outcome. In the present case this situation occurs because $w_{21}^2 = 0.03$ while $w_{22}^2 = 0.51$. Consequently, project 5 is accorded an 'efficient' status by permitting the gap between w_1^2 and w_2^2 to be (perhaps unfairly) very large. Actually, the set of multipliers which render project 5 efficient also constitute an optimal solution for project 3.

If we further constrain the model by implementing criteria importance conditions as defined in the previous section, the relative positioning of the projects changes as shown in Table 2-3.

Table 2-3. Efficiency Scores (Ranked Criteria)

Project	1	2	3	4	5	6	7	8	9	10
Score	0.71	0.72	1.00	0.60	1.00	0.80	0.62	0.63	0.50	0.35

Hence, criteria importance restrictions can have an impact on the efficiency status of the projects.

Two interesting phenomena characterize DEA problems containing ordinal data. If one examines in detail the outputs from the analysis of the example data, two observations can be made. First, it is the case that $\theta=1$ for each project (whether efficient or inefficient). This means that each project is either on the frontier proper or an extension of the frontier. Second, if one were to use the CRS rather than VRS model, it would be observed that $\sum \lambda_k = 1$ for each project. The implication would seem to be that the two models (CRS and VRS) are equivalent in the presence of ordinal data. Moreover, since $\theta = 1$ in all cases, these models are as well equivalent to the additive model of Charnes et al. (1985). The following two theorems prove these results for the general case.

Theorem 3.3

In problem (3.6'), if I_2 is non empty, $\theta = 1$ at the optimum.

Proof:

By definition, $\bar{\delta}_{ik}(\ell) = \sum_{n=\ell}^L \delta_{ik}(n)$. Thus, $\bar{\delta}_{ik}(1)=1$ for all k , and for any

ordinal input i . From the constraint set of (3.6'), if $\sum_{k=1}^N \lambda_k = 1$, then since

$$\theta \bar{\delta}_{io}(1) \geq \sum_{k=1}^N \lambda_k \bar{\delta}_{ik}(1), \text{ and since } \sum_{k=1}^N \lambda_k \bar{\delta}_{ik}(1)=1 \text{ (given that all}$$

members of $\{\bar{\delta}_{ik}(1)\}_{k=1}^N$ equal 1), it follows that $\theta \bar{\delta}_{io}(1) \geq 1$. But since $\bar{\delta}_{io}(1)=1$, then $\theta \geq 1$, meaning that at the optimum $\theta=1$.

QED

This rather unusual property of the DEA model in the presence of ordinal data is generally explainable by observing the dual form (3.5). It is noted that ε^2 plays the role of discriminating between the levels of relative importance of consecutive rank positions. If in the extreme case $\varepsilon=0$, then any one rank position becomes as important as any other. This means that regardless of the rank position occupied by a DMU "o", that position can be credited with at least as high a weight as those assumed by the peers of that DMU. Hence, every DMU will be deemed technically efficient. It is only the presence of positive gaps (defined by ε^2) between rank positions that renders a DMU inefficient via the slacks.