

Aggregating preference ranking with fuzzy Data Envelopment Analysis

Majid Zerafat Angiz L.^{a,c}, Ali Emrouznejad^{b,*}, A. Mustafa^a, A.S. Al-Eraqi^d

^a School of Mathematical Sciences, Universiti Sains Malaysia, Penang, Malaysia

^b Aston Business School, Aston University, Birmingham B4 7ET, UK

^c Department of Mathematics, Islamic Azad University, Firouzkooh, Iran

^d Department of Computer Science and Engineering, College of Engineering, University of Aden, Yemen Republic

ARTICLE INFO

Article history:

Received 27 July 2008

Received in revised form 17 March 2010

Accepted 18 March 2010

Available online 14 April 2010

Keywords:

Group decision making

Preferential voting system

Fuzzy Data Envelopment Analysis

Most preferable alternative

Aggregating preference ranking

ABSTRACT

Selecting the best alternative in a group decision making is a subject of many recent studies. The most popular method proposed for ranking the alternatives is based on the distance of each alternative to the ideal alternative. The ideal alternative may never exist; hence the ranking results are biased to the ideal point. The main aim in this study is to calculate a fuzzy ideal point that is more realistic to the crisp ideal point. On the other hand, recently Data Envelopment Analysis (DEA) is used to find the optimum weights for ranking the alternatives. This paper proposes a four stage approach based on DEA in the Fuzzy environment to aggregate preference rankings. An application of preferential voting system shows how the new model can be applied to rank a set of alternatives. Other two examples indicate the priority of the proposed method compared to the some other suggested methods.

Crown Copyright © 2010 Published by Elsevier B.V. All rights reserved.

1. Introduction

In the preference voting systems the aim is to select m alternatives from a set of n alternatives ($n > m$). Hence each expert ranks the alternatives from the most preferred (rank = 1) to the least preferred (rank = n). Obviously due to different opinions of the experts, each alternative may be placed in a different ranking position. Some studies suggest a simple aggregation method by finding the total score of each alternative as the weighted sum of the votes that each alternative receive by different experts. In this method, the best alternative is the one with the largest total score. The key issue of the preference aggregation is how to determine the weights associated with different ranking positions. Perhaps, Borda–Kendall method [15] is the most commonly used approach for determining the weights due to its computational simplicity.

In recent years, researchers have used Data Envelopment Analysis (DEA) [5,8,12,24] to rank the alternatives [1,6,7]. In the common DEA models, the objective is to determine technical efficiency while Cook and Kress [7] developed a DEA model for aggregating preferential votes. This approach proves to be effective, but sometimes there are more than one efficient DMU (Decision Making Unit), so the model is not able to discriminate the rank of some candidates. For this, Green et al. [13] suggest a method based on cross-efficiency evaluation to discriminate the candi-

dates' rank (see also Noguchi et al. [19]). Hashimoto [14] uses super efficiency DEA model [2] for ranking efficient DMUs in the Cook and Kress's model while Obata and Inguishi [20] proposed a normalized scoring vectors to distinguish the candidates' rank (see also Wang et al. [24]). Llamazares and Pena [17] reviewed some of these models (see also [13,14,20]).

Other attempts have been done using fuzzy multi-attribute [3,4,9] and fuzzy clustering methodology with ordered weighted averaging operator [4,11,16] and integrated multi-objective modeling with fuzzy multi-attribute group [18,21,23] for similar problem in group decision making.

The aim of this paper is first to obtain the ideal solution by solving a single fuzzy linear programming problem. Second, we develop a weighting method using optimization technique to find the best weights for selecting the most favorable alternative. The weights then can be used to define a social function that can fairly solve a voting problem. This paper suggests a four stage approach; Stage 1 proposes a fuzzy membership function for ranking a set of alternatives to find the ideal alternative. Stage 2 utilizes fuzzy DEA to attain the ideal solution. In Stage 3, the weights are investigated by modeling subjective evaluation with comparative judgments. Finally, Stage 4 proposes a method to aggregate the final results to single aggregated score. This score is used for ranking alternatives.

The rest of this paper is organized as follows. In Section 2 some general definitions and discussion of DEA and Fuzzy DEA is given. The proposed method is presented in Section 3. Section 4 illustrates the new method with three applications. Conclusions are given in Section 4.

* Corresponding author.

E-mail addresses: mzerafat24@yahoo.com (M. Zerafat Angiz L.), A.Emrouznejad@aston.ac.uk (A. Emrouznejad).

2. Background, DEA under fuzzy environment

2.1. Fuzzy numbers

Let us start with some definitions.

Definition 1. A fuzzy number $\tilde{x}$ is a convex normalized fuzzy set $\tilde{x}$ of the real line R such that

1. There exists exactly one $x^m \in R$ that $\mu_{\tilde{x}}(x^m) = 1$ (x^m is called the mean value of $\tilde{x}$).
2. $\mu_{\tilde{x}}(x)$ is a piecewise continuous function.

We use the definition of fuzzy numbers as in Dubois and Prade [10].

Definition 2. A fuzzy number $\tilde{x}$ is of LR-type if there exist reference functions L (for left), R (for right), and $\alpha > 0$, $\beta > 0$ such that

$$\mu_{\tilde{x}}(\bar{x}) = \begin{cases} L\left(\frac{\bar{x}-x^l}{x^m-x^l}\right) & \text{for } x^l \leq \bar{x} \leq x^m \\ R\left(\frac{x^u-\bar{x}}{x^u-x^m}\right) & \text{for } x^m \leq \bar{x} \leq x^u \end{cases}$$

x^m , called the mean value of $\tilde{x}$, is a real number, and $\alpha = x^m - x^l$ and $\beta = x^u - x^m$ are called the left and right spreads, respectively, x^l and x^u are the lower bound and the upper bound of the interval of fuzzy number as shown in Fig. 1. Symbolically, $\tilde{x}$ is denoted by $(x^m, \alpha, \beta)_{LR}$ or $(m, \alpha, \beta)_{LR}$. If $\tilde{x}$ is a symmetric triangular fuzzy number, we have $R(x) = L(x) = x$. Hence, the fuzzy number corresponding to Fig. 2 is denoted by $(m, \alpha, 0)_{LR}$.

2.2. DEA and fuzzy DEA

Data Envelopment Analysis (DEA) as introduced by Charnes et al. [5] (CCR) is a non-parametric performance assessment methodology to measure the relative efficiency of a set of homogeneous Decision Making Units (DMUs) such as bank branches, hospitals, which consume one or more inputs to produce one or more outputs.

In mathematical terms, consider a set of n DMUs, in which $x_{ij}(i = 1, 2, \dots, m)$ and $y_{rj}(r = 1, 2, \dots, s)$ are inputs and outputs of DMU_j ($j = 1, 2, \dots, n$). A standard DEA model for assessing DMU_p is formulated in Model (1) and a DEA model with non-radial measure as suggested in [13,22] is a given in Model (2). In Model (1), the $\lambda_j(j = 1, 2, \dots, n)$ are the raw weights assigned to the peer DMUs when solving the DEA model, θ_p measures the efficiency of DMU_p with input values $x_{ip}(i = 1, 2, \dots, m)$ and output values $y_{rp}(r = 1, 2, \dots, s)$. The optimal value (θ_p^*) demonstrates the relative efficiency score related to DMU_p . In this model DMU_p is efficient if $\theta_p^* = 1$. In Model (2) w_p is a free variable which measures the efficiency of DMU_p . However, in real applications, data are not often deterministic, so traditional DEA cannot be used for such problems.

Model 1: DEA, a standard CCR model	Model 2: DEA, a non-radial measure
$\min \theta_p$ $\text{s.t.} \quad \sum_{j=1}^n \lambda_j x_{ij} \leq \theta x_{ip} \quad \forall i$ $\sum_{j=1}^n \lambda_j y_{rj} \geq y_{rp} \quad \forall r$ $\lambda_j \geq 0 \quad \forall j$ $\theta \text{ free}$	$\min W_p = w_p + 1$ $\text{s.t.} \quad \sum_{j=1}^n \lambda_j x_{ij} \leq x_{ip} + w_p \quad \forall i$ $\sum_{j=1}^n \lambda_j y_{rj} \geq y_{rp} - w_p \quad \forall r$ $\lambda_j \geq 0 \quad \forall j$ $w_p \text{ free}$

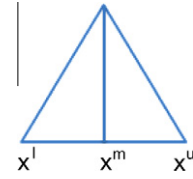


Fig. 1. A triangular fuzzy number.

For this, the non-radial DEA in Model (2) is developed to a model with fuzzy coefficients as formulated in Model (3).

Model 3: DEA, a non-radial measure with fuzzy coefficients

$$\begin{aligned} \min \quad & W_p = w_p + 1 \\ \text{s.t.} \quad & \sum_{j=1}^n \lambda_j \tilde{x}_{ij} \leq \tilde{x}_{ip} + w_p \quad \forall i \\ & \sum_{j=1}^n \lambda_j \tilde{y}_{rj} \geq \tilde{y}_{rp} - w_p \quad \forall r \\ & \lambda_j \geq 0 \quad \forall j \\ & w \text{ free} \end{aligned}$$

where “ $\sim$ ” indicates the fuzziness.

In the rest of this paper we adopt Model (3) as a base for our methodology since this model is a linear programming problem while a standard Fuzzy CCR model will lead to a nonlinear programming problem.

2.3. Cook and Kress method

Cook and Kress [7] proposed a method based on DEA to aggregate the votes from a preferential ballot. For this purpose, they used the following DEA model with no input, while outputs are number of first place votes, second place votes and so on, where $v_{ij}(i = 1, 2, \dots, m; j = 1, 2, \dots, n)$ is the number of j th place votes of the candidate i and u_j is the weight of rank j calculated by Model (4). It is clear that $u_j \geq u_{j+1}$, so the extra constraint $u_j - u_{j+1} \geq d(j, \varepsilon)$ indicates how much vote $i + 1$ is preferred to vote i . The notation $d(j, \varepsilon)$ is a function which is non-decreasing in ε and is referred to as a discrimination intensity function (see [7]). Model (4) is solved for each candidate $i = 1, 2, \dots, m$.

Model 4: Cook and Kress model

$$\begin{aligned} \max \quad & \sum_{j=1}^n u_j v_{pj} \\ \text{s.t.} \quad & \sum_{j=1}^n u_j v_{ij} \leq 1 \quad i = 1, 2, \dots, m \\ & u_j - u_{j+1} \geq d(j, \varepsilon) \quad j = 1, 2, \dots, n-1 \\ & u_n \geq d(n, \varepsilon) \end{aligned}$$

Furthermore, Cook and Kress [7] proposed a model to maximize discrimination variable ε (see also Llamazares and Pena [17]). Hashimoto [14] proposed a similar model except that they consider $d(j, \varepsilon) = \varepsilon$ for all $j = 1, 2, \dots, n$. They assumed that ε is small enough to guarantee a decreasing sequence of weights and to avoid the solution of the model depending on the discrimination intensity functions. In addition, they added the constraints $u_j - 2u_{j+1} + u_{j+2} \geq 0$, $j = 1, 2, \dots, n-2$ to Cook and Kress model in order to discriminate the efficient candidates.

Download English Version:

<https://daneshyari.com/en/article/402565>

Download Persian Version:

<https://daneshyari.com/article/402565>

[Daneshyari.com](https://daneshyari.com)