

مدل ابرکارایی AP با داده‌های ترتیبی در تحلیل پوششی داده‌های نادقیق (یک مطالعه موردی: بررسی مراکز خدمات مخابراتی کره جنوبی)

محمد خدا بخشی^{۱*}، مهدی کرمی خرم آبادی^۲، علی ثامری پور^۳

۱- دانشیار، دانشگاه شهید بهشتی، گروه ریاضی کاربردی، تهران، ایران

۲- دانشجوی کارشناسی ارشد، دانشگاه لرستان، گروه ریاضی، خرم آباد، ایران

۳- دانشیار، دانشگاه لرستان، گروه ریاضی، خرم آباد، ایران

رسید مقاله: ۲۰ دی ۱۳۹۳

پذیرش مقاله: ۵ خرداد ۱۳۹۴

چکیده

اگر چه روش‌های متفاوتی از ابرکارایی در تحلیل پوششی داده‌های کلاسیک ارایه شده است، تاکنون روش‌های ابرکارایی کلاسیک در تحلیل پوششی داده‌های ترتیبی که یکی از انواع داده‌های نادقیق می‌باشد بررسی نشده است. در این مقاله یک مدل ابرکارایی کلاسیک در تحلیل پوششی داده‌های نادقیق (IDEA) بررسی می‌شود. در واقع با فرض نادقیق بودن داده‌های ورودی و خروجی، مدل نادقیق مربوطه (AP) را تعریف نموده، برای حل مدل حاصل، معادل قطعی آن را به دست می‌آوریم. برای راحتی کار معادل قطعی مدل را که غیرخطی می‌باشد با استفاده از روش دقیق‌سازی داده‌های ترتیبی (تبدیل داده‌های ترتیبی به داده‌های دقیق)، به یک مدل بازه‌ای تبدیل می‌کنیم و سپس با اقتباس از روش دسپوتیس و اسمیرلیس مدل حاصل را به یک برنامه‌ریزی خطی تبدیل نموده، با حل آن تحت بهترین شرایط و بدترین شرایط یک جواب بهینه بازه‌ای به دست می‌آوریم که مقدار بهینه به دست آمده در آن بازه قرار دارد.

کلمات کلیدی: تحلیل پوششی داده‌ها، ابرکارایی، داده‌های ترتیبی، داده‌های بازه‌ای، داده‌های نادقیق.

۱ مقدمه

تحلیل پوششی داده‌ها (DEA)، یک ابزار تصمیم‌گیری مناسب برای ارزیابی عملکرد نسبی مجموعه‌ای از واحدهای تحت ارزیابی است که دارای ورودی و خروجی‌های مشابه می‌باشد. اولین مدل در DEA، مدل CCR است که در سال ۱۹۷۸ توسط چارنز، کوپر و رودز پیشنهاد گردید [۱]. در مدل‌های کلاسیک تحلیل پوششی داده‌ها می‌بایست، که مقدار تمامی داده‌های ورودی‌ها و خروجی‌ها شناخته شده باشد. این فرض همیشه در دنیای واقعی درست نیست. به علت عدم اطمینان در داده‌ها، برخی اوقات تحلیل پوششی داده‌ها با داده‌های نادقیق مواجه

* عهده دار مکاتبات

آدرس الکترونیکی: mkhbachshi@yahoo.com

می‌شود، مخصوصاً وقتی یک سری از واحدهای تصمیم‌گیرنده حاوی داده‌های بازه‌ای و ترتیبی باشند و به طور کلی صحبت از عدم اطمینان در داده‌ها باشد، مدل تحلیل پوششی داده‌ها به یک مدل غیرخطی تبدیل می‌شود و IDEA نامیده می‌شود.

در اغلب مدل‌های DEA، معمولاً چندین DMU کارا وجود دارد. رتبه‌بندی بین این واحدهای کارا موضوع تحقیقی جالبی است. مدل‌های مختلف و مطالعات فراوانی برای رتبه‌بندی‌های واحدهای کارا مطرح شده است از قبیل: مدل‌های اندرسن و پیترسن [۲] در سال ۱۹۹۳، دوئل و گرین [۳-۴] در سال‌های ۱۹۹۳ و ۱۹۹۴، سیفورد و ژو [۵] در سال ۱۹۹۹، ژو [۶] در سال ۲۰۰۱، لی و همکاران [۷] در سال ۲۰۰۷، خدابخشی [۸] در سال ۲۰۰۷، جهانشاهلو و همکاران [۹] در سال ۲۰۱۱، پایان و همکاران [۱۰] در سال ۲۰۱۱، خدابخشی و همکاران [۱۱] در سال ۲۰۱۲، حسین‌زاده‌لطفی و همکاران [۱۲] در سال ۲۰۱۳، رمضان و همکاران [۱۳] در سال ۲۰۱۳، جهانشاهلو و همکاران [۱۴] در سال ۲۰۱۴، و مطالعات مختلف و متعددی از این قبیل صورت گرفته، در این مقاله رتبه‌بندی را ابرکارایی می‌نامیم. ابرکارایی یک روش برای رتبه‌بندی واحدهای کارا می‌باشد که قادر است یک واحد کارا و رأسی مانند k را با مقدار کارایی بزرگ‌تر از یک مشخص کند. در کلیه مدل‌های یاد شده از داده‌های دقیق استفاده شده؛ اما از آنجایی که در اغلب مسایل عملی و کاربردی داده‌ها نادقیق و غیرقطعی است؛ لذا در سال‌های اخیر مطالعاتی در مورد ابرکارایی داده‌های نادقیق ارائه شده که از جمله آن‌ها می‌توان به تحقیقات خدابخشی و همکاران در این زمینه اشاره کرد [۱۵-۱۸]. ساختار این مقاله، بدین ترتیب است: در بخش ۲ به معرفی داده‌های نادقیق می‌پردازیم، سپس در بخش ۳ از نظر تئوریک و نظری مدل ابرکارایی AP را معرفی می‌نماییم و در ادامه در بخش ۴، مدل ابرکارایی AP را با داده‌های ترتیبی ضعیف و ترتیبی قوی توسعه داده و سرانجام یک مطالعه موردی، در خصوص داده‌های ترتیبی به طور عملی اجرا می‌نماییم، و بالاخره در آخرین قسمت این مقاله نتایج و مطالب پیشنهادی برای تحقیقات آتی پیش رو گذاشته می‌شود.

۲ داده‌های نادقیق

کوپر و همکاران [۱۹] در سال ۱۹۹۹ تکنیک تحلیل پوششی داده‌ها را در کنار داده‌های نادقیق قرار دادند. منظور از داده‌های نادقیق، داده‌هایی است که در آن‌ها ورودی‌ها و خروجی‌ها در یک بازه‌ای از اعداد تعریف شده (داده‌های بازه‌ای) و یا برای آن‌ها ترتیب در نظر گرفته می‌شود (داده‌های ترتیبی)، که ذیل آن‌ها را معرفی می‌کنیم [۲۰].

داده‌های بازه‌ای یا کراندار:

$$x_{ij}^L \leq x_{ij} \leq x_{ij}^U \quad \text{و} \quad y_{rj}^L \leq y_{rj} \leq y_{rj}^U \quad \text{برای} \quad i \in BI, r \in BO, \quad (1)$$

x_{ij}^L و y_{rj}^L کران‌های پایین و x_{ij}^U و y_{rj}^U کران‌های بالا هستند، و BI و BO به ترتیب مجموعه‌های شامل ورودی‌ها و خروجی‌های کراندار می‌باشند.

داده‌های ترتیبی ضعیف:

$$x_{ij} \leq x_{ik} \quad \text{و} \quad y_{rj} \leq y_{rk} \quad \text{برای} \quad j \neq k, i \in DI, r \in DO, \quad (2)$$

یا به طور خلاصه:

$$x_{i1} \leq x_{i2} \leq \dots \leq x_{ik} \leq \dots \leq x_{in} \quad (i \in DI), \quad (3)$$

$$y_{r1} \leq y_{r2} \leq \dots \leq y_{rk} \leq \dots \leq y_{rm} \quad (r \in DO), \quad (4)$$

که DI و DO مجموعه‌های شامل ورودی‌ها و خروجی‌هایی است که در رابطه ترتیبی ضعیف صدق می‌کند.

داده‌های ترتیبی قوی:

$$x_{i1} < x_{i2} < \dots < x_{ik} < \dots < x_{in} \quad (i \in SI), \quad (5)$$

$$y_{r1} < y_{r2} < \dots < y_{rk} < \dots < y_{rm} \quad (r \in SO), \quad (6)$$

یا به عبارتی:

$$x_{i,k} - x_{i,k-1} \geq \eta \quad (i \in SI) \quad \text{و} \quad y_{r,k} - y_{r,k-1} \geq \eta \quad (r \in SO) \quad (7)$$

که SI و SO مجموعه‌های شامل ورودی‌ها و خروجی‌هایی است که در رابطه ترتیبی قوی صدق می‌کند.

η یک متغیر مثبت است، در حالتی که $\eta = 0$ روابط فوق بیانگر روابط ترتیبی ضعیف خواهد بود. لازم به ذکر است که به جای η می‌توان عنصر غیرارشمیدسی ε را نیز جایگزین نمود.

کوپر و همکاران [۱۹]، در تحقیق‌های خود برای اولین بار، اصطلاح تحلیل پوششی داده‌های نادقیق را معرفی کردند. منظور از این اصطلاح، مدل‌هایی است که از اضافه کردن داده‌های بازه‌ای و ترتیبی، به مدل‌های کلاسیک تحلیل پوششی داده‌ها به دست آمده است.

۳ مدل ابرکارایی اندرسن - پیترسن (AP)

با به کار بردن مدل‌های DEA جهت به دست آوردن کارایی نسبی واحدهای تصمیم‌گیرنده معمولاً بیش از یک DMU کارا ارزیابی می‌شود، که رتبه‌بندی این واحدهای تصمیم‌گیرنده‌ی کارا، از اهمیت خاصی برخوردار می‌باشد. در طول سال‌های اخیر، مطالعات زیادی برای رتبه‌بندی DMU ‌های کارا صورت گرفته است. نقطه مشترک این مطالعات و تحقیقات در این است که DMU مورد ارزیابی جهت رتبه‌بندی را از مجموعه واحدهای تصمیم‌گیرنده مجموعه امکان تولید حذف نموده، و مدل DEA تعدیل یافته مورد نظر را برای محاسبه ابرکارایی DMU مورد ارزیابی با استفاده از DMU ‌های باقی مانده حل می‌کنند. معروف‌ترین و اولین این مدل‌ها، مدل پیشنهاد شده توسط اندرسن و پیترسن می‌باشد که به مدل AP معروف است که ضمن معرفی آن، این مدل با داده‌های ترتیبی نیز توسعه داده می‌شود.

مدل AP در سال ۱۹۹۳ توسط اندرسن و پیترسن [۲۱-۲۰] پیشنهاد گردید. آن‌ها جهت رتبه‌بندی، DMU_k مورد ارزیابی را از مجموعه امکان تولید حذف کردند و مجموعه امکان تولید جدیدی به صورت زیر تعریف کردند:

$$T_{AP} = \left\{ (x, y) \mid x \geq \sum_{j=1, j \neq k}^n \lambda_j x_j, y \leq \sum_{j=1, j \neq k}^n \lambda_j y_j, \lambda_j \geq 0, j = 1, \dots, n, j \neq k \right\} \quad (8)$$

در مدل AP برای محاسبه کارایی DMU_k جهت رتبه‌بندی حرکت به سوی مرز در امتداد شعاعی صورت می‌گیرد. این حرکت، مرز T_{AP} را ممکن است در نقطه‌ای قطع کند، که در این صورت مدل شدنی است. یا ممکن است T_{AP} را اصلاً قطع نکند، که در این صورت مدل نشدنی می‌گردد. یا ممکن است در فاصله بسیار دور T_{AP} را قطع کند، که در این صورت حالت ناپایداری مشاهده می‌شود. بنابر آنچه گفته شد در مدل AP، هدف پیدا کردن کم‌ترین مقدار θ_k^s است که وقتی در ورودی‌های DMU_k ضرب شود، DMU مجازی ساخته شده روی مرز T_{AP} قرار گیرد. به بیان ریاضی هدف حل مدل زیر است:

$$\begin{aligned} \text{Min } & \theta_k^s \\ \text{s.t. } & (\theta_k^s x_k, y_k) \in T_{AP} \end{aligned} \quad (9)$$

بنابر رابطه (۸)، مدل (۹) معادل مدل (۱۰) است.

$$\begin{aligned} \text{Min } & \theta_k^s \\ \text{s.t. } & \sum_{j=1, j \neq k}^n \lambda_j x_{ij} \leq \theta_k^s x_{ik}, \quad i = 1, \dots, m, \\ & \sum_{j=1, j \neq k}^n \lambda_j y_{rj} \geq y_{rk}, \quad r = 1, \dots, s, \\ & \lambda_j \geq 0, \quad j = 1, \dots, n; j \neq k, \end{aligned} \quad (10)$$

اصطلاحاً به مدل (۱۰)، فرم پوششی مدل AP (در ماهیت ورودی) گفته می‌شود.

تعریف ۱ DMU_k ابرکارا است اگر $\theta_k^{s*} > 1$ باشد و همچنین DMU_k غیرکارا است اگر $\theta_k^{s*} < 1$ باشد. دوگان مدل (۱۰) به صورت زیر است:

$$\begin{aligned} \text{Max } & \pi_k = \sum_{r=1}^s u_r y_{rk} \\ \text{s.t. } & \sum_{i=1}^m v_i x_{ik} = 1, \\ & \sum_{r=1}^s u_r y_{rj} - \sum_{i=1}^m v_i x_{ij} \leq 0, \quad j = 1, \dots, n, j \neq k, \\ & u_r \geq 0, \quad r = 1, \dots, s, \\ & v_i \geq 0, \quad i = 1, \dots, m, \end{aligned} \quad (11)$$

۴ مدل ابرکارایی AP با داده‌های ترتیبی

با کمک گرفتن از قضیه ۱ ژو [۲۰]، قضیه زیر را در ارتباط با مدل (۱۱) بیان می‌کنیم:

قضیه ۱ مدل (۱۱) را در نظر بگیرید و رابطه (۱) را به عنوان محدودیت به مدل (۱۱) اضافه کنید، فرض کنید DMU_k تحت ارزیابی باشد، جواب بهینه این مساله برای DMU_k در $x_{ik} = x_{ik}^L$ و $y_{rk} = y_{rk}^U$ و برای DMU_j ($j \neq k$) در $x_{ij} = x_{ij}^U$ و $y_{rj} = y_{rj}^L$ رخ می‌دهد.

اثبات: اثبات این قضیه مشابه قضیه ۱ ژو [۲۰]، می باشد.

۴-۱ تبدیل داده‌های ترتیبی ضعیف به داده‌های دقیق

DMU_k را در نظر بگیرید، رابطه‌های (۳) و (۴) را به عنوان محدودیت به مدل (۱۱) اضافه می کنیم و مدل به دست آمده را حل می کنیم، و یک مجموعه از جواب‌های بهینه x_{ij}^* و y_{rj}^* با مقدار تابع هدف π_k^* به دست می آوریم. داریم [۲۰]:

$$x_{i1}^* \leq x_{i2}^* \leq \dots \leq x_{i,k-1}^* \leq x_{ik}^* \leq x_{i,k+1}^* \leq \dots \leq x_{in}^* \quad (i \in DI), \quad (12)$$

$$y_{r1}^* \leq y_{r2}^* \leq \dots \leq y_{r,k-1}^* \leq y_{rk}^* \leq y_{r,k+1}^* \leq \dots \leq y_{rm}^* \quad (r \in DO), \quad (13)$$

از طرفی با توجه به خاصیت پایا بودن نسبت به تغییر واحد ρx_{ij}^* ($i \in DI$) و ρy_{rj}^* ($r \in DO$) نیز جواب بهینه برای DMU_k است، ρ که عددی ثابت و مثبت است؛ بنابراین می توانیم قرار دهیم $x_{ik}^* = y_{rk}^* = 1$ ، پس روابط (۱۲) و (۱۳) را می توان به صورت زیر نوشت:

$$0 \leq x_{i1}^* \leq x_{i2}^* \leq \dots \leq x_{i,k-1}^* \leq x_{ik}^* (=1) \leq x_{i,k+1}^* \leq \dots \leq x_{in}^* \leq M \quad (i \in DI), \quad (14)$$

$$0 \leq y_{r1}^* \leq y_{r2}^* \leq \dots \leq y_{r,k-1}^* \leq y_{rk}^* (=1) \leq y_{r,k+1}^* \leq \dots \leq y_{rm}^* \leq M \quad (r \in DO), \quad (15)$$

که M عددی به اندازه‌ی کافی بزرگ است. اکنون برای ورودی‌ها و خروجی‌هایی که در رابطه‌ی ترتیبی ضعیف صدق می کند ما بازه‌های زیر را قرار می دهیم:

$$y_{rj} \in [0, 1] \quad \text{و} \quad x_{ij} \in [0, 1] \quad \text{برای} \quad DMU_j \quad (j = 1, \dots, k-1), \quad (16)$$

$$y_{rj} \in [1, M] \quad \text{و} \quad x_{ij} \in [1, M] \quad \text{برای} \quad DMU_j \quad (j = k+1, \dots, n), \quad (17)$$

اکنون بنا بر قضیه ۱ می دانیم، که برای DMU_k اگر $i \in DI$ و $r \in DO$ باشد مقدار π_k^* بدون تغییر باقی می ماند و نیز اگر $x_{ik} = y_{rk} = 1$ ، (۱۶) و (۱۷) برقرار می شوند و از این رو برای DMU_j ($j = 1, \dots, k-1$)، $y_{rj} = 0$ (کران پایین، y_{rj}^L) و $x_{ij} = 1$ (کران بالا، x_{ij}^U) برای DMU_j ($j = k+1, \dots, n$)، $y_{rj} = 1$ (کران پایین، y_{rj}^L) و $x_{ij} = M$ (کران بالا، x_{ij}^U) به عنوان داده‌های دقیق انتخاب می شود.

۴-۲ دقیق سازی داده‌های بازه‌ای

مدل (۱۰) را در نظر بگیرید و رابطه (۱) را به عنوان محدودیت به مدل (۱۰) اضافه می کنیم، در این صورت داریم:

$$\begin{aligned}
 & \text{Min } \theta_k^s \\
 & \text{s.t.} \quad \sum_{j=1, j \neq k}^n \lambda_j [x_{ij}^L, x_{ij}^U] \leq \theta_k^s [x_{ik}^L, x_{ik}^U], \quad i = 1, \dots, m, \\
 & \quad \sum_{j=1, j \neq k}^n \lambda_j [y_{rj}^L, y_{rj}^U] \geq [y_{rk}^L, y_{rk}^U], \quad r = 1, \dots, s, \\
 & \quad \lambda_j \geq 0, \quad j = 1, \dots, n; j \neq k,
 \end{aligned} \tag{18}$$

هریک از مدل‌های DEA با داده‌های بازه‌ای یک مساله غیرخطی را در تحلیل پوششی داده‌های بازه‌ای تشکیل می‌دهد که حل آن طبق روش‌های موجود برای مسایل DEA امکان‌پذیر نیست و نیاز به استفاده از الگوریتم‌هایی خاص برای حل این گونه مدل‌هاست.

با اقتباس از روش دسپوتیس و اسمیرلیس [۲۲-۲۳]، مدل (۱۸) را حل می‌کنیم. مدل (۱۸) یک مدل بازه‌ای است پس یک مساله غیرخطی است برای تبدیل مدل غیرخطی فوق به مدل خطی تبدیل‌های زیر را انجام می‌دهیم:

$$\begin{aligned}
 & x_{ij}^L \leq x_{ij} \leq x_{ij}^U \quad 0 \leq \alpha_{ij} \leq 1 \quad x_{ij} = \alpha_{ij} x_{ij}^U + (1 - \alpha_{ij}) x_{ij}^L \Rightarrow x_{ij} = x_{ij}^L + \alpha_{ij} (x_{ij}^U - x_{ij}^L), \\
 & \quad j \neq k \quad j \neq k \\
 & x_{ik}^L \leq x_{ik} \leq x_{ik}^U \quad 0 \leq \alpha_{ik} \leq 1 \quad x_{ik} = \alpha_{ik} x_{ik}^U + (1 - \alpha_{ik}) x_{ik}^L \Rightarrow x_{ik} = x_{ik}^L + \alpha_{ik} (x_{ik}^U - x_{ik}^L), \\
 & \quad j \neq k \quad j \neq k \\
 & y_{rj}^L \leq y_{rj} \leq y_{rj}^U \quad 0 \leq \beta_{rj} \leq 1 \quad y_{rj} = \beta_{rj} y_{rj}^U + (1 - \beta_{rj}) y_{rj}^L \Rightarrow y_{rj} = y_{rj}^L + \beta_{rj} (y_{rj}^U - y_{rj}^L), \\
 & \quad j \neq k \quad j \neq k \\
 & y_{rk}^L \leq y_{rk} \leq y_{rk}^U \quad 0 \leq \beta_{rk} \leq 1 \quad y_{rk} = \beta_{rk} y_{rk}^U + (1 - \beta_{rk}) y_{rk}^L \Rightarrow y_{rk} = y_{rk}^L + \beta_{rk} (y_{rk}^U - y_{rk}^L), \\
 & \quad j \neq k \quad j \neq k
 \end{aligned} \tag{19}$$

با جای‌گزینی رابطه (۱۹) در مدل (۱۸) داریم:

$$\begin{aligned}
 & \text{Min } \theta_k^s \\
 & \text{s.t.} \quad \sum_{j=1, j \neq k}^n \lambda_j x_{ij}^L + \sum_{j=1, j \neq k}^n \lambda_j \alpha_{ij} (x_{ij}^U - x_{ij}^L) \leq \theta_k^s (x_{ik}^L + \alpha_{ik} (x_{ik}^U - x_{ik}^L)), \quad i = 1, \dots, m, \\
 & \quad \sum_{j=1, j \neq k}^n \lambda_j y_{rj}^L + \sum_{j=1, j \neq k}^n \lambda_j \beta_{rj} (y_{rj}^U - y_{rj}^L) \geq y_{rk}^L + \beta_{rk} (y_{rk}^U - y_{rk}^L), \quad r = 1, \dots, s, \\
 & \quad \lambda_j \geq 0, \quad j = 1, \dots, n; j \neq k,
 \end{aligned} \tag{20}$$

مدل (۲۰) به دلیل وجود عوامل $\lambda_j \alpha_{ij}$ و $\lambda_j \beta_{rj}$ غیرخطی است، که با تغییر متغیر $p_{ij} = \lambda_j \alpha_{ij}$ و $q_{rj} = \lambda_j \beta_{rj}$ و $j \neq k$ به برنامه‌ریزی خطی تبدیل می‌گردد، و چون $0 \leq q_{rj} \leq \lambda_j$ و $0 \leq \beta_{rj} \leq 1$ و $0 \leq p_{ij} \leq \lambda_j$ پس $0 \leq p_{ij} \leq \lambda_j$ و $0 \leq q_{rj} \leq \lambda_j$.

با اعمال این تغییر متغیرها داریم:

$$\begin{aligned}
 & \text{Min } \theta_k^s \\
 & \text{s.t.} \quad \sum_{j=1, j \neq k}^n \lambda_j x_{ij}^L + \sum_{j=1, j \neq k}^n p_{ij} (x_{ij}^U - x_{ij}^L) \leq \theta_k^s (x_{ik}^L + \alpha_{ik} (x_{ik}^U - x_{ik}^L)), \quad i = 1, \dots, m, \\
 & \quad \sum_{j=1, j \neq k}^n \lambda_j y_{rj}^L + \sum_{j=1, j \neq k}^n q_{rj} (y_{rj}^U - y_{rj}^L) \geq y_{rk}^L + \beta_{rk} (y_{rk}^U - y_{rk}^L), \quad r = 1, \dots, s, \\
 & \quad \lambda_j \geq 0, \quad j = 1, \dots, n; j \neq k, \\
 & \quad p_{ij} - \lambda_j \leq 0, \quad j = 1, \dots, n; j \neq k, \\
 & \quad q_{rj} - \lambda_j \leq 0, \quad j = 1, \dots, n; j \neq k, \\
 & \quad p_{ij} \geq 0, \quad j = 1, \dots, n; j \neq k, \\
 & \quad q_{rj} \geq 0, \quad j = 1, \dots, n; j \neq k,
 \end{aligned} \tag{21}$$

مدل (۲۱) یک مدل مینیمم سازی است در نتیجه هر چه قدر ناحیه شدنی بزرگ تر شود جواب بهینه بدتر (بزرگ تر) نمی شود؛ بنابراین برای داشتن بهترین جواب (کران پایین ابرکارایی) داریم:

$$\begin{aligned}
 & \left\{ \begin{array}{l} (x_{ik}^U - x_{ik}^L) \geq 0, \quad i = 1, \dots, m \\ 0 \leq \alpha_{ik} \leq 1, \quad i = 1, \dots, m \\ \theta_k^s \geq 0 \end{array} \right\} \alpha_{ik} = 1 \\
 & \left\{ \begin{array}{l} (y_{rk}^U - y_{rk}^L) \geq 0, \quad r = 1, \dots, s \\ 0 \leq \beta_{rk} \leq 1, \quad r = 1, \dots, s \end{array} \right\} \beta_{rk} = 0 \\
 & \left\{ \begin{array}{l} (x_{ij}^U - x_{ij}^L) \geq 0, \quad j = 1, \dots, n, j \neq k \\ 0 \leq p_{ij} \leq \lambda_j, \quad j = 1, \dots, n, j \neq k \end{array} \right\} p_{ij} = 0 \\
 & \left\{ \begin{array}{l} (y_{rj}^U - y_{rj}^L) \geq 0, \quad j = 1, \dots, n, j \neq k \\ 0 \leq q_{rj} \leq \lambda_j, \quad j = 1, \dots, n, j \neq k \end{array} \right\} q_{rj} = \lambda_j
 \end{aligned} \tag{22}$$

با اعمال روابط (۲۲) در مدل (۲۱) مدل زیر به دست می آید:

$$\begin{aligned}
 & I_k^* = \text{Min } \theta_k^s \\
 & \text{s.t.} \quad \sum_{j=1, j \neq k}^n \lambda_j x_{ij}^L \leq \theta_k^s x_{ik}^U, \quad i = 1, \dots, m, \\
 & \quad \sum_{j=1, j \neq k}^n \lambda_j y_{rj}^U \geq y_{rk}^L, \quad r = 1, \dots, s, \\
 & \quad \lambda_j \geq 0, \quad j = 1, \dots, n; j \neq k,
 \end{aligned} \tag{23}$$

به طور مشابه برای داشتن بدترین جواب (کران بالای ابرکارایی) باید پارامترها را چنان اختیار کنیم که ناحیه شدنی تا حد امکان کوچک گردد؛ لذا با اعمال شرایط زیر نتیجه مطلوب به دست می آید.

$$\left\{ \begin{array}{l} (x_{ik}^U - x_{ik}^L) \geq 0, \quad i = 1, \dots, m \\ 0 \leq \alpha_{ik} \leq 1 \quad i = 1, \dots, m \end{array} \right\} \alpha_{ik} = 0$$

$$\left\{ \begin{array}{l} (y_{rk}^U - y_{rk}^L) \geq 0, \quad r = 1, \dots, s \\ 0 \leq \beta_{rk} \leq 1 \quad r = 1, \dots, s \end{array} \right\} \beta_{rk} = 1 \quad (24)$$

$$\left\{ \begin{array}{l} (x_{ij}^U - x_{ij}^L) \geq 0, \quad j = 1, \dots, n, j \neq k \\ 0 \leq p_{ij} \leq \lambda_j \quad j = 1, \dots, n, j \neq k \end{array} \right\} p_{ij} = \lambda_j$$

$$\left\{ \begin{array}{l} (y_{rj}^U - y_{rj}^L) \geq 0, \quad j = 1, \dots, n, j \neq k \\ 0 \leq q_{rj} \leq \lambda_j \quad j = 1, \dots, n, j \neq k \end{array} \right\} q_{rj} = 0$$

با اعمال رابطه (24) در مدل (21) مدل زیر به دست می‌آید:

$$u_k^* = \min \theta_k^s$$

$$s.t. \quad \sum_{j=1, j \neq k}^n \lambda_j x_{ij}^U \leq \theta_k^s x_{ik}^L, \quad i = 1, \dots, m,$$

$$\sum_{j=1, j \neq k}^n \lambda_j y_{rj}^L \geq y_{rk}^U, \quad r = 1, \dots, s,$$

$$\lambda_j \geq 0, \quad j = 1, \dots, n; j \neq k, \quad (25)$$

بنابراین هر DMU که با مدل (18) ارزیابی می‌گردد، مقدار ابرکارایی آن در یک بازه قرار می‌گیرد، که کران پایین ابرکارایی؛ یعنی u_k^* از حل مدل (23) و کران بالای ابرکارایی؛ یعنی u_k^* از حل مدل (25) به دست می‌آید. در قسمت 4-1 داده‌های ترتیبی ضعیف را به داده‌های بازه‌ای تبدیل کردیم، طبق روش دقیق‌سازی داده‌های بازه‌ای که در قسمت 4-2 بیان شد می‌توان کران‌های ابرکارایی را برای داده‌های ترتیبی ضعیف به صورت زیر به دست آورد.

کران پایین ابرکارایی نیز وقتی به دست می‌آید که:

$$x_{ij} = 0 \text{ برای } DMU_j (j = 1, \dots, k-1) \text{ و } x_{ij} = 1 \text{ برای } DMU_j (j = k+1, \dots, n), \quad (26)$$

$$y_{rj} = 1 \text{ برای } DMU_j (j = 1, \dots, k-1) \text{ و } y_{rj} = M \text{ برای } DMU_j (j = k+1, \dots, n), \quad (27)$$

کران بالای ابرکارایی وقتی حاصل می‌شود که:

$$x_{ij} = 1 \text{ برای } DMU_j (j = 1, \dots, k-1) \text{ و } x_{ij} = M \text{ برای } DMU_j (j = k+1, \dots, n), \quad (28)$$

$$y_{rj} = 1 \text{ برای } DMU_j (j = 1, \dots, k-1) \text{ و } y_{rj} = M \text{ برای } DMU_j (j = k+1, \dots, n), \quad (29)$$

4-3 مثال کاربردی

داده‌های ورودی و خروجی 8 مرکز خدماتی مخبرات کره جنوبی را که در جدول (1) ارائه شده است در نظر بگیرید. این داده‌ها مربوط به سال 1996 می‌باشد و از پایگاه‌های اطلاعاتی به دست آمده است [24].

شاخص‌های ورودی عبارت است از:

۱. نیروی کار (x_1): داده‌ها متناظر آن در جدول بیانگر تعداد کارمندان می‌باشد. این تعداد شامل نیروهای متخصص برای انجام اموری مانند مسولیت تکنولوژی جدید و تحقیق روی خدمات ماهواره‌ای نمی‌باشد.
۲. هزینه عملیات (x_2): هزینه‌ی مربوط به ارایه خدمات تلفن‌ها و پیجرهای موبایل است. این هزینه شامل هزینه سود، استهلاک، هزینه سرمایه‌گذاری تجهیزات و هزینه‌های کار نمی‌باشد.
۳. سطح مدیریت برای تسهیلات و مشتریان (x_3): این ورودی، تجربه تیم پرسنل را در تسهیل‌سازی و برخورد با مشتریان نشان می‌دهد. بخش مدیریت این عامل را برای تمام شعبه‌ها سالی یک‌بار ارزیابی می‌کند. و نتایج را مبنای افزایش حقوق‌ها، ترفیع‌ها و پاداش‌ها قرار می‌دهد. داده‌های ستون سوم تنها بیانگر وجود روابط ترتیبی بین مولفه‌های این ورودی است. برای مثال، Kwangju بالاترین رتبه را در بین سایر شعب دارد و Jeju کم‌ترین رتبه را دارد. در واقع داده‌های این ستون بیانگر رابطه‌ی ترتیبی زیر می‌باشد:

$$x_{34} \geq x_{35} \geq x_{33} \geq x_{37} \geq x_{31} \geq x_{36} \geq x_{32} \geq x_{38}$$

شاخص‌های خروجی نیز عبارتند از:

۱. درآمد (y_1): بیانگر مبلغ قبض‌های ارایه خدمات تلفن‌ها و پیجرهای موبایل است.
۲. میزان خرابی وسایل (y_2): تعداد خرابی وسایل را به منظور تعویض نشان می‌دهد.
۳. میزان تکمیل مکالمه (y_3): این شاخص نشان دهنده‌ی نسبت تعداد تماس‌های موفق بر تعداد کل تماس‌هاست. میزان تکمیل مکالمه هر شعبه به طور زیادی وابسته به مشخصه‌های جغرافیایی منطقه است. تاثیر چشمگیر عواملی مانند کوهستان‌ها، زیرگذرها، ساختمان‌های بزرگ و ... روی این خروجی باعث می‌شود که برای میزان تکمیل مکالمه مقدار دقیقی حاصل نشود و این مقدار برای هر شعبه‌ای در محدوده‌های کراندار واقع شود.

جدول ۱. داده‌های مراکز خدماتی مخابرات کره جنوبی

DMU_j	name	(x_1)	(x_2)	(x_3)	(y_1)	(y_2)	(y_3)
۱	Seoul	۱۲۴	۱۸/۲۲	۴	۲۵/۵۳	۸۹/۸	[۸۰,۸۵]
۲	Pusan	۹۵	۹/۲۳	۲	۱۸/۴۳	۹۹/۶	[۸۵,۹۰]
۳	Taegu	۹۲	۸/۰۷	۶	۱۰/۲۹	۸۷	[۷۵,۸۰]
۴	Kwangju	۶۱	۵/۶۲	۸	۸/۳۲	۹۹/۴	۱۰۰
۵	Taejun	۶۳	۵/۳۳	۷	۷/۰۴	۹۶/۴	[۷۰,۷۵]
۶	jeonju	۵۰	۳/۵۳	۳	۶/۴۲	۸۶	[۹۰,۹۵]
۷	Kangnung	۴۰	۳/۵	۵	۲/۲	۷۱	[۸۰,۸۵]
۸	jeju	۱۶	۱/۱۷	۱	۲/۸۷	۹۸	[۹۵,۱۰۰]

معمولا عدد M را برابر با تعداد واحد ها در نظر می گیرند. در این مثال $M = 8$.

همان گونه که بیان شد رابطه بین داده‌های ورودی سوم جدول (۱) یک رابطه تریبی و هم چنین داده‌های خروجی سوم جدول (۱) به صورت داده بازه‌ای است. با توجه به روابط (۲۶) و (۲۷) فرم دقیق این داده‌ها برای به دست آوردن کران پایین ابرکارایی در جدول (۲) آمده است.

جدول ۲. فرم دقیق داده‌های X_3 و Y_3 برای محاسبه کران پایین ابرکارایی

DMU_j	$\frac{DMU_1}{x_3 \quad y_3}$	$\frac{DMU_2}{x_3 \quad y_3}$	$\frac{DMU_3}{x_3 \quad y_3}$	$\frac{DMU_4}{x_3 \quad y_3}$	$\frac{DMU_5}{x_3 \quad y_3}$	$\frac{DMU_6}{x_3 \quad y_3}$	$\frac{DMU_7}{x_3 \quad y_3}$	$\frac{DMU_8}{x_3 \quad y_3}$
۱	۱ ۸۰	۱ ۸۵	۰ ۸۵	۰ ۸۵	۰ ۸۵	۱ ۸۵	۰ ۸۵	۱ ۸۵
۲	۰ ۹۰	۱ ۸۵	۰ ۹۰	۰ ۹۰	۰ ۹۰	۰ ۹۰	۰ ۹۰	۱ ۹۰
۳	۱ ۸۰	۱ ۸۰	۱ ۷۵	۰ ۸۰	۰ ۸۰	۱ ۸۰	۱ ۸۰	۱ ۸۰
۴	۱ ۱۰۰	۱ ۱۰۰	۱ ۱۰۰	۱ ۱۰۰	۱ ۱۰۰	۱ ۱۰۰	۱ ۱۰۰	۱ ۱۰۰
۵	۱ ۷۵	۱ ۷۵	۱ ۷۵	۰ ۷۵	۱ ۷۰	۱ ۷۰	۱ ۷۵	۱ ۷۵
۶	۰ ۹۵	۱ ۹۵	۰ ۹۵	۰ ۹۵	۰ ۹۵	۱ ۹۰	۰ ۹۵	۱ ۹۵
۷	۱ ۸۰	۱ ۸۵	۰ ۸۵	۰ ۸۵	۰ ۸۵	۱ ۸۵	۱ ۸۰	۱ ۸۵
۸	۰ ۱۰۰	۰ ۱۰۰	۰ ۱۰۰	۰ ۱۰۰	۰ ۱۰۰	۰ ۱۰۰	۰ ۱۰۰	۱ ۹۵

به طور مشابه، با توجه به روابط (۲۸) و (۲۹) فرم دقیق داده‌های X_3 و Y_3 برای محاسبه کران بالای ابرکارایی در جدول (۳) آمده است.

جدول ۳. فرم دقیق داده‌های X_3 و Y_3 برای محاسبه کران بالای ابرکارایی

DMU_j	$\frac{DMU_1}{x_3 \quad y_3}$	$\frac{DMU_2}{x_3 \quad y_3}$	$\frac{DMU_3}{x_3 \quad y_3}$	$\frac{DMU_4}{x_3 \quad y_3}$	$\frac{DMU_5}{x_3 \quad y_3}$	$\frac{DMU_6}{x_3 \quad y_3}$	$\frac{DMU_7}{x_3 \quad y_3}$	$\frac{DMU_8}{x_3 \quad y_3}$
۱	۱ ۸۵	۸ ۸۰	۱ ۸۰	۱ ۸۰	۱ ۸۰	۸ ۸۰	۱ ۸۰	۸ ۸۰
۲	۱ ۸۵	۱ ۹۰	۱ ۸۵	۱ ۸۵	۱ ۸۵	۱ ۸۵	۱ ۸۵	۸ ۸۵
۳	۸ ۷۵	۸ ۷۵	۱ ۸۰	۱ ۷۵	۱ ۷۵	۱ ۷۵	۸ ۷۵	۸ ۷۵
۴	۸ ۱۰۰	۸ ۱۰۰	۸ ۱۰۰	۱ ۱۰۰	۸ ۱۰۰	۸ ۱۰۰	۸ ۱۰۰	۸ ۱۰۰
۵	۸ ۷۰	۸ ۷۰	۸ ۷۰	۱ ۷۰	۱ ۷۵	۸ ۷۰	۸ ۷۰	۸ ۷۰
۶	۱ ۹۰	۸ ۹۰	۱ ۹۰	۱ ۹۰	۱ ۹۰	۱ ۹۵	۱ ۹۰	۸ ۹۰
۷	۱ ۸۰	۸ ۸۰	۱ ۸۰	۱ ۸۰	۱ ۸۰	۸ ۸۰	۱ ۸۵	۸ ۸۰
۸	۱ ۹۵	۱ ۹۵	۱ ۹۵	۱ ۹۵	۱ ۹۵	۱ ۹۵	۱ ۹۵	۱ ۱۰۰

نتایج ارزیابی کران‌های پایین (l_k^*) و بالای (u_k^*) ابرکارایی (به ترتیب با استفاده از مدل‌های (۲۳) و (۲۵)) در جدول (۴) آمده است.

جدول ۴. نتایج عددی

DMU_j	l_k^*	u_k^*
۱	۱/۰۶۱۳	۱/۳۸۵۲
۲	۱/۰۳۱۹	۵/۸۱۵۰
۳	۰/۵۹۴۱	۰/۹۱۶۰
۴	۰/۷۱۵۸	۱/۰۸۸۸
۵	۰/۵۹۹۶	۰/۹۹۹۱
۶	۰/۷۴۱۴	۱/۲۷۳۲
۷	۰/۳۲۰۰	۰/۸۹۴۷
۸	۳/۵۴۴۱	۸/۰۰۰

۴-۴ تبدیل داده‌های ترتیبی قوی به داده‌های دقیق

با کمک گرفتن از قضیه ۱ ژو [۲۵]، قضیه زیر را در ارتباط با مدل (۱۱) بیان می‌کنیم:

قضیه ۲ برای هر η ، اضافه کردن رابطه (۷) به مدل (۱۱) معادل است با اضافه کردن رابطه

$$x_{i,k} - x_{i,k-1} \geq \varepsilon \text{ و } y_{r,k} - y_{r,k-1} \geq \varepsilon \text{ به مدل (۱۱) در جایی که } \varepsilon \approx 0.$$

اثبات: اثبات این قضیه مشابه قضیه ۱ ژو [۲۵]، می‌باشد.

روند دقیق‌سازی داده‌های ترتیبی قوی مشابه داده‌های ترتیبی ضعیف است. DMU_k را در نظر بگیرید،

رابطه‌های (۵) و (۶) را به عنوان محدودیت به مدل (۱۱) اضافه می‌کنیم و مدل به‌دست آمده را حل می‌کنیم، و

یک مجموعه از جواب‌های بهینه x_{ij}^* و y_{rj}^* به‌دست می‌آوریم. داریم [۲۴]:

$$x_{i1}^* \leq x_{i2}^* \leq \dots \leq x_{i,k-1}^* \leq x_{ik}^* \leq x_{i,k+1}^* \leq \dots \leq x_{in}^* \quad (i \in DI), \quad (30)$$

$$y_{r1}^* \leq y_{r2}^* \leq \dots \leq y_{r,k-1}^* \leq y_{rk}^* \leq y_{r,k+1}^* \leq \dots \leq y_{rm}^* \quad (r \in DO), \quad (31)$$

از طرفی با توجه به خاصیت پایا بودن نسبت به تغییر واحد ρx_{ij}^* ($i \in DI$) و ρy_{rj}^* ($r \in DO$) نیز جواب بهینه

برای DMU_k است، ρ که عددی ثابت و مثبت است؛ بنابراین می‌توانیم قرار دهیم $x_{ik}^* = y_{rk}^* = 1$. پس روابط

(۳۰) و (۳۱) را می‌توان به صورت زیر نوشت:

$$x_{i1}^* \leq x_{i2}^* \leq \dots \leq x_{i,k-1}^* \leq x_{ik}^* (=1) \leq x_{i,k+1}^* \leq \dots \leq x_{in}^* \quad (i \in DI), \quad (32)$$

$$y_{r1}^* < y_{r2}^* < \dots \leq y_{r,k-1}^* \leq y_{rk}^* (=1) \leq y_{r,k+1}^* \leq \dots \leq y_{rm}^* \quad (r \in SO), \quad (33)$$

به عبارت دیگر:

$$y_{r,j+1}^* - y_{rj}^* \geq \eta \text{ و } x_{i,j+1}^* - x_{ij}^* \geq \eta, \quad (34)$$

با تلفیق رابطه‌های (۳۲) و (۳۳) و (۳۴) می‌توان نوشت:

$$x_{ij}^* (=y_{rj}^*) \leq 1 - (k-j)\eta \text{ برای } DMU_j \quad (j=1, \dots, k-1), \quad (35)$$

$$x_{ij}^* (= y_{rj}^*) \geq 1 + (j - k)\eta \quad \text{برای } DMU_j \quad (j = k + 1, \dots, n), \quad (36)$$

از طرفی با توجه به قضیه ۲، همواره یک عدد به اندازه کافی کوچک مانند $\varepsilon \approx 0$ موجود است به طوری که:

$$\varepsilon \leq x_{ij} = y_{rj} \leq 1 - (k - j)\eta \quad \text{برای } DMU_j \quad (j = 1, \dots, k - 1), \quad (37)$$

$$1 + (j - k)\eta \leq x_{ij} = y_{rj} \leq M \quad \text{برای } DMU_j \quad (j = k + 1, \dots, n), \quad (38)$$

که M یک عدد به اندازه کافی بزرگ است. در این صورت عوامل ورودی و خروجی بهینه را می‌توان به فرم داده‌های بازه‌ای بیان کرد. برای سایر داده‌های ورودی و خروجی نیز این مطلب برقرار است؛ بنابراین برای حاصل شدن کران بالای ابرکارایی قرار می‌دهیم:

$$x_{ij} = 1 - (k - j)\eta \quad \text{برای } DMU_j \quad (j = 1, \dots, k - 1) \text{ و } x_{ij} \geq M \quad \text{برای } DMU_j \quad (j = k + 1, \dots, n), \quad (39)$$

$$y_{rj} \leq \varepsilon \quad \text{برای } DMU_j \quad (j = 1, \dots, k - 1) \text{ و } y_{rj} = 1 + (j - k)\eta \quad \text{برای } DMU_j \quad (j = k + 1, \dots, n), \quad (40)$$

برای محاسبه‌ی کران پایین ابرکارایی نیز قرار می‌دهیم:

$$x_{ij} \leq \varepsilon \quad \text{برای } DMU_j \quad (j = 1, \dots, k - 1) \text{ و } x_{ij} = 1 + (j - k)\eta \quad \text{برای } DMU_j \quad (j = k + 1, \dots, n), \quad (41)$$

$$y_{rj} = 1 - (k - j)\eta \quad \text{برای } DMU_j \quad (j = 1, \dots, k - 1) \text{ و } y_{rj} \geq M \quad \text{برای } DMU_j \quad (j = k + 1, \dots, n), \quad (42)$$

مقدار ابرکارایی مدل‌های DEA با روابط ترتیبی قوی وابسته به مقدار η است؛ یعنی برای محاسبه‌ی مقدار ابرکارایی لازم است که مقدار η را تعیین کنیم. در این صورت عوامل ورودی و خروجی به صورت داده‌های دقیق برآورد می‌شود و مساله غیرخطی IDEA به یک مساله برنامه‌ریزی خطی قابل حل تبدیل می‌گردد.

برای تعیین η مدل (۱۱) را با داده‌های ورودی و خروجی ترتیبی قوی در نظر بگیرید، مدل حاصل به

صورت زیر است:

$$\begin{aligned} \text{Max} \quad & \sum_{r=1}^s u_r y_{rk} \\ \text{s.t.} \quad & \sum_{i=1}^m v_i x_{ik} = 1, \end{aligned} \quad (43)$$

$$\sum_{r=1}^s u_r y_{rj} - \sum_{i=1}^m v_i x_{ij} \leq 0, \quad j = 1, \dots, n; j \neq k,$$

$$u_r \geq 0, \quad r = 1, \dots, s,$$

$$v_i \geq 0, \quad i = 1, \dots, m,$$

$$x_{ij} - x_{i,j+1} \geq \eta, x_{i,j+1} - x_{i,j+2} \geq \eta, \dots, x_{in} \geq \eta,$$

$$y_{rj} - y_{r,j+1} \geq \eta, y_{r,j+1} - y_{r,j+2} \geq \eta, \dots, y_{rm} \geq \eta,$$

مدل فوق یک مدل غیرخطی است که در آن $x_{ij} = \max_j \{x_{ij}\}$ و $y_{ij} = \min_j \{y_{ij}\}$ با به کار بردن

$$Y_r = \mu_r y_r; \quad \forall r$$

$$X_i = v_i x_i; \quad \forall i$$

تکنیک تغییر متغیر:

مدل غیرخطی مذکور به مدل خطی زیر تبدیل می گردد:

$$\begin{aligned} \text{Max} \quad & \sum_{r=1}^s Y_{rk} \\ \text{s.t.} \quad & \sum_{i=1}^m X_{ik} = 1, \\ & \sum_{r=1}^s Y_{rj} - \sum_{i=1}^m X_{ij} \leq 0, \quad j = 1, \dots, n; j \neq k, \\ & X_{ij} - X_{i,j+1} \geq \eta, X_{i,j+1} - X_{i,j+2} \geq \eta, \dots, X_{in} \geq \eta, \\ & Y_{rj} - Y_{r,j+1} \geq \eta, Y_{r,j+1} - Y_{r,j+2} \geq \eta, \dots, Y_{rm} \geq \eta, \end{aligned} \quad (44)$$

که $y_{ij} = \min_j \{y_{ij}\}$ و $x_{ij} = \max_j \{x_{ij}\}$

با اقتباس از روش کوپر و همکارانش [۲۴] برای تعیین η مدل زیر را معرفی می کنیم که با محاسبه مقدار بهینه

η_j^* برای هر واحد، کمترین مقدار آن‌ها را به عنوان η در نظر می گیریم، مدل پیشنهادی به صورت زیر است:

$$\begin{aligned} \eta_k^* = \text{Max} \quad & \eta \\ \text{s.t.} \quad & \sum_{i=1}^m X_{ik} = 1, \\ & \sum_{r=1}^s Y_{rj} - \sum_{i=1}^m X_{ij} \leq 0, \quad j = 1, \dots, n; j \neq k, \\ & X_{ij} - X_{i,j+1} \geq \eta, X_{i,j+1} - X_{i,j+2} \geq \eta, \dots, X_{in} \geq \eta, \\ & Y_{rj} - Y_{r,j+1} \geq \eta, Y_{r,j+1} - Y_{r,j+2} \geq \eta, \dots, Y_{rm} \geq \eta, \end{aligned} \quad (45)$$

η مورد نظر نیز برابر است با $\eta = \min_j \{\eta_j^*\}$.

۴-۵ مثال کاربردی

مثال کاربردی بخش ۳-۳ را در نظر بگیرید. فرض کنید داده‌های ورودی سوم در جدول (۱) بیانگر رابطه‌ی ترتیبی قوی زیر باشد.

$$X_{34} > X_{35} > X_{33} > X_{37} > X_{31} > X_{36} > X_{32} > X_{38}$$

این رابطه را می توان به صورت زیر بازنویسی کرد:

$$\begin{aligned} X_{34} - X_{35} &\geq \eta, X_{35} - X_{33} \geq \eta, X_{33} - X_{37} \geq \eta, X_{37} - X_{31} \geq \eta, \\ X_{31} - X_{36} &\geq \eta, X_{36} - X_{32} \geq \eta, X_{32} - X_{38} \geq \eta, X_{38} \geq \eta, \end{aligned}$$

چون داده‌های خروجی به صورت ترتیبی قوی نمی باشد؛ بنابراین مقدار η را می توان از حل مدل زیر محاسبه کرد:

خدا، بخشی و بهکاران، مدل ابرکارایی AP با داده‌های تریبی در تحلیل پوششی داده‌های نادرین (یک مطالعه موردی: بررسی مراکز خدمات بهداشتی کره جنوبی)

$$\eta_k^* = \text{Max } \eta$$

$$s.t. \quad \sum_{i=1}^m X_{ik} = 1,$$

$$\sum_{r=1}^s Y_{rj} - \sum_{i=1}^m X_{ij} \leq 0, \quad j = 1, \dots, n; j \neq k,$$

$$X_{34} - X_{35} \geq \eta, X_{35} - X_{33} \geq \eta, X_{33} - X_{37} \geq \eta, X_{37} - X_{31} \geq \eta,$$

$$X_{31} - X_{36} \geq \eta, X_{36} - X_{32} \geq \eta, X_{32} - X_{38} \geq \eta, X_{38} \geq \eta,$$

مقادیر η_j^* برای هر DMU در جدول (۵) آمده است.

جدول ۵. مقادیر بهینه η_j^*

DMU_j	η_j^*
۱	۰/۲۵۰
۲	۰/۵
۳	۰/۱۶۷
۴	۰/۱۲۵
۵	۰/۱۴۳
۶	۰/۳۳۳
۷	۰/۲۰۰
۸	۱

کمترین مقدار η_j^* به DMU_4 اختصاص دارد که برابر ۰/۱۲۵ می‌باشد. از این رو $\eta = 0/125$.

با توجه به رابطه (۳۹)، فرم دقیق داده‌های X_3 برای محاسبه کران بالای ابرکارایی در جدول (۶) آمده است.

جدول ۶. مقادیر دقیق (X_3) برای محاسبه کران بالای ابرکارایی

DMU_j	DMU_1^*	DMU_2	DMU_3	DMU_4	DMU_5	DMU_6	DMU_7	DMU_8
۱	۱	۹	۰/۷۵	۰/۵	۰/۶۲۵	۸	۰/۸۷۵	۱۰
۲	۰/۷۵	۱	۰/۵	۰/۲۵	۰/۳۷۵	۰/۸۷۵	۰/۶۲۵	۸
۳	۹	۱۱	۱	۰/۷۵	۰/۸۷۵	۱۰	۸	۱۲
۴	۱۱	۱۳	۹	۱	۸	۱۲	۱۰	۱۴
۵	۱۰	۱۲	۸	۰/۸۷۵	۱	۱۱	۹	۱۳
۶	۰/۸۷۵	۸	۰/۶۲۵	۰/۳۷۵	۰/۵	۱	۰/۷۵	۹
۷	۸	۱۰	۰/۸۷۵	۰/۶۲۵	۰/۷۵	۹	۱	۱۱
۸	۰/۶۲۵	۰/۸۷۵	۰/۳۷۵	۰/۱۲۵	۰/۲۵	۰/۷۵	۰/۵	۱

*داده‌ها در این ستون نشان داده شده مقادیر دقیق X_3 برای هر DMU_j ($j = 1, \dots, 8$) هنگامی که

DMU_1 تحت ارزیابی باشد.

فرم دقیق داده‌های بردار بازه ای γ_3 برای محاسبه کران بالای ابرکارایی نیز همانند جدول (۳) می‌باشد. نتایج حاصل از محاسبه کران بالای ابرکارایی (با استفاده از مدل (۲۵)) در جدول (۷) آمده است:

جدول ۷. نتایج عددی

DMU_j	u_k^*
۱	۱/۰۶۱۳
۲	۵/۶۱۸۹
۳	۰/۶۰۲۵
۴	۰/۷۱۶۶
۵	۰/۵۹۹۶
۶	۰/۸۴۴۳
۷	۰/۴۴۷۴
۸	۹/۴۱۱۸

اگر $\eta = ۰/۰۶۲۵$ را بگیریم، در این صورت با توجه به رابطه (۳۹) فرم دقیق داده‌های X_3 برای محاسبه کران بالای ابرکارایی در جدول (۸) آمده است.

جدول ۸. مقادیر دقیق (X_3) برای محاسبه کران بالای ابرکارایی

DMU_j	DMU_1	DMU_2	DMU_3	DMU_4	DMU_5	DMU_6	DMU_7	DMU_8
۱	۱	۹	۰/۸۷۵	۰/۷۵	۰/۸۱۲۵	۸	۰/۹۳۷۵	۱۰
۲	۰/۸۷۵	۱	۰/۷۵	۰/۶۲۵	۰/۶۸۷۵	۰/۹۳۷۵	۰/۸۱۲۵	۸
۳	۹	۱۱	۱	۰/۸۷۵	۰/۹۳۷۵	۱۰	۸	۱۲
۴	۱۱	۱۳	۹	۱	۸	۱۲	۱۰	۱۴
۵	۱۰	۱۲	۸	۰/۹۳۷۵	۱	۱۱	۹	۱۳
۶	۰/۹۳۷۵	۸	۰/۸۱۲۵	۰/۶۸۷۵	۷۵	۱	۰/۸۷۵	۹
۷	۸	۱۰	۰/۹۳۷۵	۰/۸۱۲۵	۰/۸۷۵	۹	۱	۱۱
۸	۰/۸۱۲۵	۰/۹۳۷۵	۰/۶۸۷۵	۰/۵۶۲۵	۰/۶۲۵	۰/۸۷۵	۰/۷۵	۱

فرم دقیق داده‌های بردار بازه ای γ_3 برای محاسبه کران بالای ابرکارایی نیز همانند جدول (۳) می‌باشد. نتایج حاصل از محاسبه کران بالای ابرکارایی (با استفاده از مدل (۲۵)) در جدول (۹) آمده است:

خدا، بخشی و به کاران، مدل ابرکارایی AP با داده‌های تریبی در تحلیل پوششی داده‌های نادقیق (یک مطالعه موردی: بررسی مراکز خدمات خبراتی کره جنوبی)

جدول ۹. نتایج عددی

DMU_j	u_k^*
۱	۱/۲۱۲۱
۲	۶/۰۲۰۳
۳	۰/۶۴۵۹
۴	۰/۷۱۶۶
۵	۰/۶۲۹۰
۶	۰/۹۰۹۶
۷	۰/۶۷۱۱
۸	۹/۴۱۱۸

مقایسه کران بالای ابرکارایی در جدول‌های (۷) و (۹) نشان می‌دهد که هر چه مقدار η بزرگ‌تر باشد مقدار ابرکارایی بدتر (u_k^* تغییری نمی‌کند یا کم‌تر می‌شود) خواهد شد. با اختیار کردن $\varepsilon = 7 \times 10^{-5}$ و با توجه به رابطه (۴۱) فرم دقیق داده‌های x_p برای محاسبه کران پایین ابرکارایی مطابق جدول (۱۰) ارایه شده است.

جدول ۱۰. مقادیر دقیق (x_p) برای محاسبه کران پایین ابرکارایی

DMU_j	DMU_1	DMU_2	DMU_3	DMU_4	DMU_5	DMU_6	DMU_7	DMU_8
۱	۱	۱/۲۵	۰/۰۰۰۰۶	۰/۰۰۰۰۴	۰/۰۰۰۰۵	۱/۱۲۵	۰/۰۰۰۰۷	۱/۳۷۵
۲	۰/۰۰۰۰۶	۱	۰/۰۰۰۰۴	۰/۰۰۰۰۲	۰/۰۰۰۰۳	۰/۰۰۰۰۷	۰/۰۰۰۰۵	۱/۱۲۵
۳	۱/۲۵	۱/۵	۱	۰/۰۰۰۰۶	۰/۰۰۰۰۷	۱/۳۷۵	۱/۱۲۵	۱/۶۲۵
۴	۱/۵	۱/۷۵	۱/۲۵	۱	۱/۱۲۵	۱/۶۲۵	۱/۳۷۵	۱/۸۷۵
۵	۱/۳۷۵	۱/۶۲۵	۱/۱۲۵	۰/۰۰۰۰۷	۱	۱/۵	۱/۲۵	۱/۷۵
۶	۰/۰۰۰۰۷	۱/۱۲۵	۰/۰۰۰۰۵	۰/۰۰۰۰۳	۰/۰۰۰۰۴	۱	۰/۰۰۰۰۶	۱/۲۵
۷	۱/۱۲۵	۱/۳۷۵	۰/۰۰۰۰۷	۰/۰۰۰۰۵	۰/۰۰۰۰۶	۱/۲۵	۱	۱/۵
۸	۰/۰۰۰۰۵	۰/۰۰۰۰۷	۰/۰۰۰۰۳	۰/۰۰۰۰۱	۰/۰۰۰۰۲	۰/۰۰۰۰۶	۰/۰۰۰۰۴	۱

فرم دقیق داده‌های بردار بازه‌ای y_p برای محاسبه کران پایین ابرکارایی نیز همانند جدول (۲) می‌باشد. نتایج حاصل از محاسبه کران پایین ابرکارایی (با استفاده از مدل (۲۳)) در جدول (۱۱) ارایه شده است:

جدول ۱۱. نتایج عددی

DMU_j	I_k^*
۱	۱/۰۶۱۳
۲	۱/۰۳۱۹
۳	۰/۵۹۴۱
۴	۰/۷۱۵۸
۵	۰/۵۹۹۶
۶	۰/۷۴۱۴
۷	۰/۳۲۰۰
۸	۳/۵۴۴۱

۴-۶ رده بندی ابر کارایی

فرض کنید I_k^* کران پایین ابر کارایی و u_k^* کران بالای ابر کارایی در مسایل $IDEA$ باشد، در این صورت بر مبنای مقادیر ابر کارایی، رده بندی زیر از DMU ها انجام می شود:

$$1- DMU_k \text{ ابر کارا است هرگاه: } I_k^* > 1$$

$$2- DMU_k \text{ ناکارا است هرگاه: } u_k^* < 1$$

۵ نتیجه گیری و پیشنهادات

در این مقاله مدل ابر کارایی اندرسن - پیترسن در حالت بازده به مقیاس ثابت با داده های ترتیبی در تحلیل پوششی داده های نادقیق توسعه داده شده است. و چگونگی تبدیل مساله اندازه گیری ابر کارایی در $IDEA$ به یک مساله خطی قابل حل در DEA شرح داده شده است و توانستیم جواب ابر کارایی را در مسایل تحلیل پوششی داده های ترتیبی تعیین نماییم، برای تحقیقات آتی پیشنهاد می گردد مدل ابر کارایی اندرسن - پیترسن نیز با فرض داده های فازی توسعه داده شود.

منابع

[۲۳] نورا، ع.ع.، یحیی پور م.ت.، (۱۳۸۴). رتبه بندی تحلیل پوششی داده های فازی بر مبنای داده های آماری و فاصله اطمینان آماری. مجله تحقیقات در عملیات در کاربردهای آن، (۷)، ۴۰-۴۴.

- [1] Charnes, A., Cooper, W. W., Rhodes, E., (1978). Measuring the efficiency of decision making units. European Journal of Operational Research, 2(6): 429-444.
- [2] Andersen, P., Petersen, N. C., (1993). A procedure for ranking efficient units in data envelopment analysis. Management Science, 39(10): 1261-1264.
- [3] Doyle, J., Green, R., (1993). Data Envelopment Analysis and Multiple Criteria Decision Making. Omega, 21(6): 713-715.
- [4] Doyle, J., Green, R., (1994). Efficiency and Cross- Efficiency in DEA: Derivations, meanings and uses. Journal of the Operational Research Society, 45(5): 567-578.
- [5] Seiford, L. M., Zhu, J., (1999). Infeasibility of Super Efficiency Data Envelopment Analysis Models. INFORS, 37(2): 174-187.

- [6] Zhu, J., (2001). Super Efficiency and DEA Sensitivity Analysis. *European Journal of Operational Research*, 129(2): 443-455.
- [7] Li, S., Jahanshahloo, G. R., Khodabakhshi, M., (2007). A super-efficiency model for ranking efficient units in data envelopment analysis. *Applied Mathematics and Computation*, 184(2): 638-648.
- [8] Khodabakhshi, M., (2007). A super-efficiency model based on improved outputs in data envelopment analysis. *Applied Mathematics and Computation*, 184(2): 695-703.
- [9] Jahanshahloo, G. R., Khodabakhshi, M., Hosseinzadeh Lotfi, F., Moazami Goudarzi, MR., (2011). A cross-efficiency model based on super-efficiency for ranking units through the TOPSIS approach and its extension to the interval case. *Mathematical and Computer Modelling*, 53(9-10): 1946-1955.
- [10] Payan, A., Hosseinzadeh Lotfi, F., Noora, AA., Khodabakhshi, M., (2011). A modified common set of weights method to complete ranking DMUs. *Interpretation*, 5(7): 1143-1153.
- [11] Khodabakhshi, M., Aryavash, K., (2012). Ranking all units in data envelopment analysis. *Applied Mathematics Letters*, 25(12): 2066-2070.
- [12] Hosseinzadeh Lotfi, F., Jahanshahloo, G. R., Khodabakhshi, M., Rostamy-Malkhlifeh, M., Moghaddas, Z., Vaez-Ghasemi, M., (2013). A review of ranking models in data envelopment analysis. *Journal of Applied Mathematics*, Volume 2013 (2013), Article ID 492421, 20 pages.
- [13] Ramazani-Tarkhorani, S., Khodabakhshi, M., Mehrabian, S., Nuri-Bahmani, F., (2013). Ranking decision-making units using common weights in DEA. *Applied Mathematical Modelling*. In press.
- [14] Jahanshahloo, H., Hosseinzadeh Lotfi, F., Khodabakhshi, M., (2014). Modified nonradial super efficiency models. *Journal of Applied Mathematics*, Volume 2014 (2014), Article ID 919547, 5 pages.
- [15] Khodabakhshi, M., Asgharian, M., Gregoriou, G. N., (2010). An input-oriented super-efficiency measure in stochastic data envelopment analysis: Evaluating chief executive officers of US public banks and thrifts. *Expert Systems with Applications*, 37(3): 2092-2097.
- [16] Khodabakhshi, M., (2010). An output oriented super-efficiency measure in stochastic data envelopment analysis: Considering Iranian electricity distribution companies. *Computers & Industrial Engineering*, 58(4): 663-671.
- [17] Khodabakhshi, M., (2011). Super-efficiency in stochastic data envelopment analysis: An input relaxation approach. *Journal of computational and applied mathematics*, 235(16): 4576-4588.
- [18] Khodabakhshi, M., Aryavash, K., (2014). Ranking units with fuzzy data in DEA. *Data Envelopment Analysis and Decision Science*, Volume 2014 (2014), Article ID dea-00058, 10 Pages.
- [19] Cooper, W. W., Park, K. S., Yu, G., (1999). IDEA and AR-IDEA: Models for dealing with imprecise data in DEA. *Management Science*, 45(4): 597-607.
- [20] Zhu, J., (2003). Imprecise data envelopment analysis (IDEA): A review and improvement with an application. *European Journal of Operational Research*, 144(3): 513-529.
- [21] Adle, N., Friedman, L., Sinuany-Stern, Z., (2002). Review of ranking methods in the data envelopment analysis context. *European Journal of Operational Research*, 140(2): 249-265.
- [22] Despotis, D. K., Smirlise, Y., (2002). Data envelopment analysis with imprecise data. *European Journal of Operational Research*, 140(1): 24-36.
- [24] Cooper, W. W., Park, K. S., Yu, G., (2001). An illustrative application of IDEA (imprecise data envelopment analysis) to a Korean mobile. to a Korean mobile telecommunication company. *Operations Research*, 49(6): 807-820.
- [25] Zhu, J., (2003). Efficiency evaluation with strong ordinal input and output measures. *European Journal of Operational Research*, 146(3): 477-485.