

# روشی برای رتبه بندی نمرات کارایی با استفاده از بوت استرپ

سعید عبادی\*

دانشگاه آزاد اسلامی، واحد اردبیل، گروه ریاضی، اردبیل، ایران

رسید مقاله: ۸۹/۱۰/۲۸

پذیرش مقاله: ۹۰/۲/۱

## چکیده

تعدادی از واحدهای تصمیم گیرنده که در مدل های DEA کارا ارزیابی می شوند دارای نمره ی کارایی یک هستند، پس هیچ تمایزی بین آن ها - از لحاظ تئوری - وجود ندارد. اغلب تصمیم گیرنده ها در صدد رتبه بندی کاملی از DMU ها هستند، تا بتوانند عملکرد آن ها را بهتر ارزیابی کنند و در جهت بهبود بیشتر پیش روند. تا به حال روش های مختلفی، برای رتبه بندی واحدهای کارا و ناکارا ارایه شده و مقالات زیادی در این زمینه منتشر یافته است. که بر هر یک از این روش ها ایرادهایی وارد است و آن ها قادر به رتبه بندی تمامی DMU ها نیستند. روش بوت استرپ مطرح شده در این مقاله می تواند رتبه بندی تمامی واحدها را انجام دهد. روند بوت استرپ به عوض تحلیل تئوری از قدرت محاسباتی استفاده می کند. و برای بوت استرپ نمرات کارایی به طور ضمنی از برنامه ریزی خطی استفاده می کند. و به همین دلیل هیچکدام از مشکلات روش های قبلی را ندارد.

**کلمات کلیدی:** تحلیل پوششی داده ها، واحدهای تصمیم گیرنده، مجموعه ی امکان تولید، بوت استرپ.

## ۱ مقدمه

تحلیل پوششی داده ها (DEA) شامل تکنیک ها و روش هایی برای ارزیابی کارایی و یا سنجش بهره وری واحدهای تصمیم گیرنده (DMU) است. به منظور ارزیابی علمی عملکرد واحدها و بخش ها، فعالیت های علمی زیادی صورت گرفته است. با توجه به تمایل رو به رشد تصمیم گیرنده ها و عوامل اجرایی به رتبه بندی کلی DMU ها، یک مجموعه ی جامع از روش های رتبه بندی و نتایج حاصل از مقایسه ی آن ها، می تواند نقش بسزایی در سرعت و دقت انتخاب مدیران و مجریان داشته باشد. به طوری که هر تصمیم گیرنده می تواند با توجه به فرایند جاری و واحدهای تحت ارزیابی، روشی از روش های رتبه بندی را طوری انتخاب کند که واحدهای مد نظر به صورت جامع تر و با دقت بیشتری متمایز شوند.

\*عهده دار مکاتبات

تعریف رابطه‌ی عملکرد با روابط تاثیر گذار به ساخت تابعی با عنوان تابع تولید منجر شده که از ترکیب ورودی‌ها بیشترین خروجی ممکن را تولید می‌کند. واضح است که به دست آوردن تابعی که در تعریف بالا بگنجد، کار دشوار و در بسیاری از موارد غیر ممکن است. پس برای به دست آوردن تابع، آن را به روش‌های مختلف تقریب زدند. از جمله‌ی این روش‌ها می‌توان به روش‌های پارامتری و غیر پارامتری اشاره کرد. با پیشرفت تکنولوژی، روش‌های پارامتری در برخورد با مسایل موفق عمل نکرد. برای رفع مشکلات ناشی از روش‌های پارامتری، فارل در سال ۱۹۵۷ برای نخستین بار روش‌های غیر پارامتری را ابداع کرد.

DEA، در واقع تعمیم کار فارل در ابداع اولین روش غیر پارامتری است. فارل، با استفاده از ورودی‌ها و خروجی‌های واحد‌های تصمیم‌گیرنده و با استفاده از اصول حاکم بر آن‌ها، مجموعه‌ای با عنوان مجموعه‌ی امکان تولید را ارایه و قسمتی از مرز آن را تقریبی از تابع تولید خواند. این مرز را مرز کارا نیز می‌نامند و واحدهای تصمیم‌گیرنده‌ای که روی این مرز قرار می‌گیرند، کارا ارزیابی می‌شوند. از آنجائی که DEA تکنیک ارزیابی کارایی نسبی واحدهای تصمیم‌گیرنده است، حداقل یکی از واحدها روی مرز و بقیه‌ی واحدها در زیر آن قرار دارند. نام تحلیل پوششی داده‌ها، از ویژگی پوششی بودن منشا گرفته است. این روش‌ها در مقایسه با بقیه روش‌های عملی قبلی، دارای مزیت‌هایی است که در ادامه به ذکر آن می‌پردازیم.

در روش‌های DEA بر خلاف بعضی از روش‌های عددی، مشخص بودن وزن‌ها از قبل و تخصیص آن‌ها به ورودی و خروجی داده‌ها لازم نیست، همچنین این روش‌ها، نیازی به اشکال تابعی از قبل مشخص شده (مانند روش رگرسیون‌های آماری) و یا شکل صریح تابع تولید (مانند برخی از روش‌های پارامتری) ندارند.

تحلیل پوششی داده‌ها، با استفاده از تکنیک برنامه‌ریزی ریاضی، می‌تواند تعداد زیادی از متغیرها و روابط را در برگیرد و مشکلات روش‌هایی را که در استفاده از ورودی‌ها و خروجی‌ها با محدودیت مواجه‌اند را ندارد. به علاوه بدنه‌ی وسیعی از فرضیه و تکنیک‌های در دسترس از برنامه‌ریزی ریاضی، می‌تواند به عنوان تحلیل‌ها و تعبیر راهنما و همچنین محاسبات تاثیر گذار استفاده شود.

همچنین DEA، فرصت‌های زیادی برای همکاری میان تحلیل‌گر و تصمیم‌گیرنده ایجاد می‌کند. این همکاری‌ها می‌تواند در راستای انتخاب ورودی و خروجی واحدهای تحت ارزیابی و چگونگی عملکرد و الگویابی نسبت به مرز کارا باشد.

## ۲ مدل‌های DEA و رتبه‌بندی

مجموعه‌ای به صورت  $\{X \text{ ورودی } X \text{ خروجی } Y \text{ را تولید کند: } (X, Y)\}$  که در آن  $X$  بردار ورودی و  $Y$  بردار خروجی است را مجموعه‌ی امکان تولید (PPS) می‌نامند. این تعریف با توجه به نوع تکنولوژی تولید متفاوت، PPS‌های مختلف را تولید می‌کنند.

اولین PPS که به مجموعه‌ی امکان تولید CCR معروف است در سال ۱۹۷۸ توسط چارلز و همکارانش [3] به صورت زیر تعریف شد:

Min  $\theta$

$$\begin{aligned}
 s.t. \quad & \sum_{j=1}^n \lambda_j x_{ij} \leq \theta x_{io} \quad i = 1, 2, \dots, m, \\
 & \sum_{j=1}^n \lambda_j y_{rj} \geq y_{ro} \quad r = 1, 2, \dots, s, \\
 & \lambda_j \geq 0 \quad j = 1, 2, \dots, n.
 \end{aligned} \tag{۱}$$

مدل دیگری که برای تحلیل و بررسی مفهوم کارایی در تحلیل پوششی داده ها، مورد استفاده قرار می گیرد، مدل جمعی است که توسط کوپر و همکاران [4] معرفی گردیده است. این مدل هم دارای ماهیت ورودی و هم دارای ماهیت خروجی است. در این مدل ماکزیمم فاصله ی واحد تصمیم گیرنده از مرز کارایی سنجیده می شود.

$$\begin{aligned}
 Max \quad & z = \sum_{i=1}^m s_i^- + \sum_{r=1}^s s_r^+ \\
 s.t. \quad & \sum_{j=1}^n \lambda_j x_{ij} + s_i^- = x_{io} \quad i = 1, 2, \dots, m, \\
 & \sum_{j=1}^n \lambda_j y_{rj} - s_r^+ = y_{ro} \quad r = 1, 2, \dots, s, \\
 & \sum_{j=1}^n \lambda_j = 1, \\
 & \lambda_j \geq 0 \quad j = 1, 2, \dots, n, \\
 & s_i^- \geq 0 \quad i = 1, 2, \dots, m, \\
 & s_r^+ \geq 0 \quad r = 1, 2, \dots, s.
 \end{aligned} \tag{۲}$$

این مدل، مدل جمعی BCC [2] است که با حذف قید  $\sum_{j=1}^n \lambda_j = 1$  به مدل جمعی CCR تبدیل می شود. در مدل جمعی واحد تحت ارزیابی کاراست اگر و فقط اگر  $z^* = 0$  که  $z^*$  مقدار بهینه ی تابع هدف را نشان می دهد.

مدل متغیر کمکی تعدیل یافته (SA) که توسط سوییشی و همکاران [18] در سال ۱۹۹۹ ارایه شد مبتنی بر اعمال محدودیت های وزنی بر روی بردارهای  $u$  و  $v$  می باشد:

$$\begin{aligned}
 Max \quad & \sum_{r=1}^s u_r y_{ro}, \\
 s.t. \quad & \sum_{i=1}^m v_i x_{io} = 1, \\
 & \sum_{r=1}^s u_r y_{rj} - \sum_{i=1}^m v_i x_{ij} \leq 0, \quad j = 1, 2, \dots, n, \\
 & u_r \geq 1/((m+s)R_r^+), \quad r = 1, 2, \dots, s, \\
 & v_i \geq 1/((m+s)R_i^-), \quad i = 1, 2, \dots, m.
 \end{aligned} \tag{۳}$$

که در آن

$$R_r^+ = \text{Max}\{y_{rj}\} \quad 1 \leq j \leq n, \quad R_i^- = \text{Max}\{x_{ij}\} \quad 1 \leq j \leq n$$

اندرسون و پیترسون [1] در سال ۱۹۹۳ مدل ابرکارایی را برای رتبه‌بندی واحدهای کارایی راسی معرفی کردند، که به مدل AP مشهور است. آن‌ها جهت تعیین رتبه‌ی واحد تصمیم‌گیرنده‌ی (0)، آن را از مجموعه‌ی امکان تولید حذف کردند و مدل را برای باقیمانده‌ی DMU اجرا نمودند. مدل پیشنهادی آن‌ها برای رتبه‌بندی  $DMU_0$  چنین است:

$$\begin{aligned} \text{Max } z &= \sum_{r=1}^s u_r y_{ro} \\ \text{s.t. } \quad &\sum_{i=1}^m v_i x_{io} = 1, \\ &\sum_{r=1}^s u_r y_{rj} - \sum_{i=1}^m v_i x_{ij} \leq 0, \quad j = 1, 2, \dots, n, \quad j \neq 0, \\ &u_r \geq \varepsilon, \quad r = 1, 2, \dots, s, \\ &v_i \geq \varepsilon, \quad i = 1, 2, \dots, m. \end{aligned} \quad (4)$$

مدل AP دارای مشکلات اساسی زیر است:

الف. اندرسون و پیترسون در مدل مربوط، مقدار تابع هدف را به عنوان رتبه‌ی همه‌ی واحدها در نظر گرفتند، بر خلاف اینکه هر واحد با وزن‌های مختلفی ارزیابی شده است. مقدار تابع هدف این مدل در واقع نسبت  $\text{Max}$  کارایی هر واحد را به نزدیکترین واحد حقیقی روی مرز جدید به دست می‌دهد.

ب. مدل AP در ماهیت ورودی برای DMU‌های با ورودی‌های خاص، ممکن است نشدنی گردد. به این معنی که اگر در ارزیابی  $DMU_0$ ، یکی از ورودی‌های آن صفر باشد مثلاً  $(x_{k0} = 0)$  و DMU‌ی دیگری وجود نداشته باشد که ورودی  $k$ ام آن صفر باشد، آنگاه مدل (۲،۲)، ممکن است نشدنی گردد. البته مدل AP در ماهیت خروجی همواره شدنی است.

ج. در بعضی از موارد ممکن است، تغییرات کوچکی در داده‌ها، تغییرات زیادی را در  $\theta^*$  حاصل کند، البته این مشکل در ماهیت خروجی مدل رخ نمی‌دهد.

د. مدل ابرکارایی AP، در رتبه‌بندی واحدهای غیر راسی ناتوان است.

مدل رتبه‌بندی AP از اولین روش‌های رتبه‌بندی ارائه شده است که در این گروه جای می‌گیرد. این مدل در دو ماهیت ورودی و خروجی طراحی شده است که بر حسب نیاز به کار گرفته می‌شود. مدل AP در ماهیت ورودی، ممکن است برای بعضی از داده‌ها نشدنی شود، همچنین برای داده‌های نزدیک به صفر امکان ناپایداری در مدل وجود دارد، البته این مشکلات در ماهیت خروجی وجود ندارد. سادگی و راحتی استفاده از این مدل باعث شده که علی‌رغم ضعف‌ها و کمبودهایی که در این روش وجود دارد، کاربرد وسیعی در همه‌ی زمینه‌ها داشته باشد.

مدل AP برای واحدهای تصمیم گیرنده ای که ورودی آن ها نزدیک به صفر است، رتبه بندی منطقی به دست نمی دهد و  $\theta^*$  حاصله ممکن است بسیار بزرگ باشد. به این دلیل مدل AP در ماهیت ورودی و برای ساختار خاصی از داده ها نا پایدار است.

سوییشی [18] در سال ۱۹۹۹ کران های خاصی را بر روی وزن های ورودی و خروجی، در مدل ابر کارایی تعریف کرد. او همانند مدل های پیشین، واحد تصمیم گیرنده ی تحت ارزیابی را از مجموعه حذف نمود و مدل متغیر کمکی را برای بقیه ی واحد ها در نظر گرفت. مدل پیشنهادی سوییشی برای رتبه بندی واحدهای کارا به صورت زیر است:

$$\begin{aligned} \text{Min } & \theta - \left( \sum_{i=1}^m \frac{s_i^-}{R_i^-} + \sum_{r=1}^s \frac{s_r^+}{R_r^+} \right) / (m + s) \\ \text{s.t. } & \sum_{\substack{j=1 \\ j \neq o}}^n \lambda_j x_{ij} + s_i^- = \theta x_{io}, \quad i = 1, 2, \dots, m, \\ & \sum_{\substack{j=1 \\ j \neq o}}^n \lambda_j y_{rj} - s_r^+ = y_{ro}, \quad r = 1, 2, \dots, s, \\ & \lambda_j \geq 0, \quad j = 1, 2, \dots, n, j \neq o, \\ & s_r^+ \geq 0, \quad r = 1, 2, \dots, s, \\ & s_i^- \geq 0, \quad i = 1, 2, \dots, m. \\ & R_r^+ = \text{Max}\{y_{rj}\}, \quad 1 \leq j \leq n, \quad R_i^- = \text{Max}\{x_{ij}\}, \quad 1 \leq j \leq n \end{aligned} \quad (5)$$

مقدار تابع هدف مساله ی ۵ را با  $\delta_0^*$  نشان می دهند و آن را به عنوان یک عدد نشانه در نظر می گیرند. در این مساله مشکل نشدنی بودن که در مدل AP مطرح بود، همچنان باقی است.

پس از آن سوییشی و همکاران، با اعمال محدودیت های وزنی بر روی وزن های ورودی و خروجی مدل ابر کارایی AP، بعضی از مشکلات ناشی از صفر شدن وزن ها را بر طرف کردند. این مدل که تحت عنوان مدل متغیر کمکی تعدیل یافته ارایه شد، این مدل هم مانند مدل قبل برای بعضی از داده ها نشدنی است. همچنین سوییشی برای رفع مشکلات ناشی از ناپایداری، عدد نشانه ی AIN را پیشنهاد کرد. این عدد همواره در بازه ی [1,2] قرار دارد.

روش بهبود خروجی نیز یک روش شعاعی است، که در جهت افزایش خروجی ها به کار می رود. این مدل بسیار شبیه به مدل AP در ماهیت خروجی است ولی در استفاده از ورودی ها انعطاف پذیری بیشتری دارد. این مدل برای انواع داده ها، شدنی و کراندار است. انعطاف پذیری مدل مذکور در استفاده از ورودی ها، توجه بسیاری از مدیران اجرایی را به خود جلب کرده است.

با توجه به مشکلات اساسی مدل AP، محرابیان و همکاران [16] در سال ۱۹۹۹ مدل دیگری را ( که به مدل MAJ معروف است ) جهت رتبه بندی DMU های کارا مطرح کردند. در مدل AP حرکت به سوی مرز در

امتداد شعاعی صورت می‌گرفت که ممکن بود سطح پوششی PPS را قطع نکند. در این صورت مساله جواب شدنی ندارد. و یا ممکن بود در فاصله‌ی بسیار دور، سطح پوششی PPS را قطع نماید. که در این صورت نیز مساله ناپایدار می‌شد. در مدل ابر کارایی غیر شعاعی MAJ حرکت به سوی مرز در امتداد محورهای ورودی، در ماهیت ورودی و با قدم‌های مساوی انجام می‌شود. که مشکل ناپایداری را رفع می‌کند ولی مشکل نشدنی بودن برای بعضی داده‌ها، در این مدل همچنان وجود دارد.

برای برطرف کردن این مشکل، ساعتی و همکاران [17] در سال ۱۹۹۹ مدلی را پیشنهاد کردند که در آن کاهش ورودی‌ها همزمان با افزایش خروجی‌ها - با اندازه‌ی مساوی - DMU تحت ارزیابی را روی مرز کارایی تصویر کند. این مدل همواره شدنی است.

در سال ۲۰۰۵ جهانشاهلو و همکاران [11] نشان دادند که تکنیک استفاده شده در مدل MAJ در جهت پایداری این مدل تحت تغییر واحد باعث بروز مشکلاتی در نتایج رتبه‌بندی می‌شود. اگر واحدهای کارا و در نتیجه مجموعه‌ی امکان تولید جدید در مدل مذکور ثابت بماند، اینطور انتظار می‌رود که نتایج حاصل از مدل رتبه‌بندی نیز ثابت باشد، اما در عمل اینطور نیست.

مشاهدات نشان می‌دهد که تغییر در ورودی بعضی از واحدهای ناکارا - جائیکه بقیه‌ی داده‌ها ثابت بماند - نتایج حاصل از رتبه‌بندی را تغییر می‌دهد. همچنین نرمالیزه کردن ورودی‌ها در مدل MAJ، ممکن است باعث تغییر نتایج رتبه‌بندی در اثر تغییرات واحدهای ناکارا - اگر چه واحدهای کارا و PPS تغییر نکنند - شود. جهانشاهلو و همکاران [10] با شناخت و بررسی علل این مشکلات، تکنیک نرمال کردن را به گونه‌ای تغییر دادند که این ناپایداری را برطرف نمایند. آن‌ها در نرمال کردن، به جای استفاده از بیشینه‌ی ورودی‌ها در هر مولفه تحت عنوان  $R_i$ ، از Max ورودی‌های واحدهای کارا، تحت عنوان  $M_i$  استفاده کردند.

پس از آن با نسبت دادن وزن‌هایی به گام‌های ورودی و خروجی، مدل جامع JHF را ارایه کردند. این وزن‌ها بر اساس اهمیت گام‌های ورودی و خروجی، توسط مدیر تعیین می‌شود. تعیین وزن‌ها در نتایج رتبه‌بندی بسیار موثر است، با اختصاص دادن بردار وزنی خاص به ورودی‌ها و خروجی‌ها، این مدل به مدل‌های MAJ و MAJ اصلاح شده تبدیل می‌شود. مدل مذکور برای همه‌ی داده‌ها شدنی است. تابع هدف این مدل به تابع هزینه یا قیمت موسوم است و در فرایند اجرای مدل Min می‌شود.

پروفسور لی و همکاران [14] مدل ابر کارایی غیر شعاعی LJK را در سال ۲۰۰۷ مطرح کردند، این مدل نیز برای بعضی از داده‌ها نشدنی است. شرایط شدنی بودن مدل همانند مدل MAJ است. گام‌های ورودی در این روش بر خلاف مدل MAJ متغیر است، این انعطاف‌پذیری در تغییر ورودی‌ها، در بسیاری از موارد اقتصادی و مدیریتی حائز اهمیت است.

جهانشاهلو و همکاران [12] در سال ۲۰۰۴ مدل جدیدی را برای رتبه‌بندی واحدهای کارایی با استفاده از نرم یک ارایه نمودند که همانند بعضی از مدل‌های پیشین برای رتبه‌بندی واحدهای کارایی غیر راسی پیشنهادی ندارد. در این مدل - همانند مدل ابر کارایی - واحد تحت ارزیابی از مجموعه‌ی امکان تولید حذف شده و مجموع انحرافات به صورت زیر کمینه می‌شود:

$$\begin{aligned}
 \text{Min } \Gamma_c^o(x, y) &= \sum_{i=1}^m |x_i - x_{io}| + \sum_{r=1}^s |y_r - y_{ro}| \\
 \text{s.t. } \sum_{\substack{j=1 \\ j \neq o}}^n \lambda_j x_{ij} &\leq x_i, & i = 1, 2, \dots, m, \\
 \sum_{\substack{j=1 \\ j \neq o}}^n \lambda_j y_{rj} &\geq y_r, & r = 1, 2, \dots, s, \\
 \lambda_j &\geq 0, & j = 1, 2, \dots, n, \quad j \neq o, \\
 y_r &\geq 0, & r = 1, 2, \dots, s, \\
 x_i &\geq 0, & i = 1, 2, \dots, m.
 \end{aligned} \tag{۶}$$

بار دیگر جهانشاهلو و همکاران [9] سعی کردند با استفاده از نرم چیشف، قدر مطلق ماکزیمم انحرافات مولفه ای را Min کنند. و مدل پیشنهادی را به صورت زیر فرموله کردند:

$$\begin{aligned}
 \text{Min } \|A - \bar{A}\|_{\infty} \\
 \text{s.t. } \sum_{\substack{j=1 \\ j \neq o}}^n \lambda_j x_{ij} &\leq \bar{x}_i, & i = 1, 2, \dots, m, \\
 \sum_{\substack{j=1 \\ j \neq o}}^n \lambda_j y_{rj} &\geq \bar{y}_r, & r = 1, 2, \dots, s, \\
 \lambda_j &\geq 0, & j = 1, 2, \dots, n, \quad j \neq o.
 \end{aligned} \tag{۷}$$

یکی از مشکلات مدل ابر کارایی، به دست آمدن مجموعه وزن های مجزا برای DMU های متفاوت است، که باعث می شود مقایسه ی درستی بین واحدهای تصمیم گیرنده انجام نگیرد و همچنین از نظر اقتصادی تعریف قابل قبولی ندارد. یکی از راههایی که برای رفع این مشکل پیشنهاد شده، روش رتبه بندی با مجموعه وزن های مشترک (CSW) است.

حسین زاده لطفی و همکاران [5] با توجه به این مساله که ناحیه ی شدنی همه ی این مسایل یکی است، مدلی را به صورت زیر ارائه کردند که در آن همه ی DMU توسط یک (v,u) منحصر به فرد ارزیابی شوند.

$$\begin{aligned}
 & \text{Max} \left\{ \frac{\sum_{r=1}^s u_r y_{r1}}{\sum_{i=1}^m v_i x_{i1}}, \dots, \frac{\sum_{r=1}^s u_r y_{rm}}{\sum_{i=1}^m v_i x_{im}} \right\} \\
 & \text{s.t.} \quad \frac{\sum_{r=1}^s u_r y_{rj}}{\sum_{i=1}^m v_i x_{ij}} \leq 1, \quad j = 1, 2, \dots, n, \\
 & \quad u_r \geq \varepsilon, \quad r = 1, 2, \dots, s, \\
 & \quad v_i \geq \varepsilon, \quad i = 1, 2, \dots, m.
 \end{aligned} \tag{8}$$

در واقع در این مدل واحدها در بهترین شرایط ارزیابی نمی شوند و برای هر واحد یک مدل حل نمی شود. در سال ۲۰۰۵ جهانشاهلو و همکاران [11] روش دیگری برای یافتن مجموعه وزن های مشترک، در مدل های DEA به صورت زیر پیشنهاد کردند:

$$\begin{aligned}
 & \text{Max} \left\{ \frac{\sum_{r=1}^s u_r y_{r1} + u_o}{\sum_{i=1}^m v_i x_{i1}}, \dots, \frac{\sum_{r=1}^s u_r y_{rm} + u_o}{\sum_{i=1}^m v_i x_{im}} \right\} \\
 & \text{s.t.} \quad \frac{\sum_{r=1}^s u_r y_{rj} + u_o}{\sum_{i=1}^m v_i x_{ij}} \leq 1, \quad j = 1, 2, \dots, n, \\
 & \quad u_r \geq \varepsilon, \quad r = 1, 2, \dots, s, \\
 & \quad v_i \geq \varepsilon, \quad i = 1, 2, \dots, m.
 \end{aligned} \tag{9}$$

در این مدل، روش تبدیل برنامه ریزی کسری چند هدفه به یک برنامه ریزی غیر خطی، متفاوت است. یکی از جدیدترین مدل های رتبه بندی در غالب رتبه بندی با مجموعه وزن های مشترک CWA است که با وزن دار کردن مجموعه ی داده ها و تعیین فاصله از خط الگو، واحدها را از هم متمایز می کند. برای روشن تر شدن بیشتر مطلب به [15] رجوع کنید.

از مزیت های کاربردی این مدل – که در مسایل مشخص می شود – می توان موارد ذیل را ذکر کرد:

۱. این مدل واحدهایی با ورودی و خروجی های خیلی بزرگ را نیز رتبه بندی می کند.
  ۲. این مدل واحدهای تصمیم گیرنده ای که اندیس ورودی و خروجی آن ها از  $n/2$  تجاوز می کند را نیز رتبه بندی می کند.
- از مشکلات این روش به دست آمدن جواب های بهینه ی چند گانه است که منجر به نتایج رتبه بندی متفاوت می شود.

جهانشاهلو و همکاران [8] در سال ۲۰۰۸ با استفاده از روش شبیه سازی مونت کارلو، روشی را برای رتبه بندی

همه ی واحدهای کارا ارایه نمودند بر خلاف مشکلاتی که ممکن است در استفاده از این روش با آن مواجه شویم، به خاطر اینکه این روش در رتبه بندی واحدهای کارای غیر راسی موفق است، ارزشمند است. جهانشاهلو و همکاران [6] در سال ۲۰۰۶، روش رتبه بندی مبتنی بر مرز کاملاً ناکارا را ارایه کردند، در این روش با تعریف مرز کاملاً ناکارا، به جای ارزیابی هر واحد تصمیم گیرنده نسبت به مرز کارا، فاصله ی آن را از مرز کاملاً ناکارا محاسبه می شود. این فاصله هم از طریق شعاعی و هم غیر شعاعی قابل محاسبه است. در مقایسه ی این روش با مدل های رتبه بندی که تا به حال ارایه شده است، به نتایج زیر می رسیم:

- این مدل مشکلی در جهت جواب های بهینه ی چند گانه ندارد.
- مدل پیشنهادی همواره شدنی و پایدار است.
- این روش قادر به رتبه بندی واحدهای کارای غیر راسی است.
- در این روش واحدهایی با بدترین عملکرد ( که در بعضی موارد مدیریتی و اقتصادی حائز اهمیت است ) مشخص می شوند. دیگر برتری های این مدل، پایداری تحت انتقال است که به خاطر همین ویژگی، این مدل را می توان برای داده های صفر و منفی هم بکار برد.
- مشکلی که این روش با آن مواجه است، واحدهایی هستند که روی فصل مشترک مرز کارا و ناکارا قرار دارند. این واحدها دارای بهترین ورودی ( خروجی ) در بعضی مولفه ها و بدترین ورودی ( خروجی ) در مولفه های دیگر هستند. این مدل نتایج قابل تحلیلی برای این مدل ها به دست نمی دهد.
- اکثر روش هایی که مورد مطالعه قرار گرفته اند، مانند روش های ذکر شده و روش هایی که جهانشاهلو [7] و جی [13] اخیراً ارایه کردند مبتنی بر خاصیت غالب و مغلوب بودن واحدها عمل می کنند. و آن ها را در دو گروه کارا و ناکارا تقسیم بندی و سپس به رتبه بندی واحدهای کارا می پردازند.
- برخی از این روش ها، از تکنیک حذف DMU از مجموعه ی امکان تولید (PPS) و بررسی تغییرات دیگر DMU ها در PPS جدید استفاده می کنند. اینگونه روش ها به روش های ابر کارایی معروف است.
- در این دسته، بعضی از روش ها، Min فاصله ی DMU ی حذف شده را از PPS جدید، معیاری برای رتبه ی آن واحد قرار می دهند. فاصله ی مذکور به روش های مختلف محاسبه می شود:

۱. روش های شعاعی: در این روش ها، فاصله ی DMU ی حذف شده، در امتداد شعاعی از مرز کارای PPS جدید، محاسبه می شود.

۲. روش های غیر شعاعی: در این روش ها، حرکت در جهت موازی محورهای ورودی و خروجی و به سوی مرز کارا است.

۳. استفاده از توابع نرم: در اینجا فاصله ی مذکور با استفاده از نرمهای مختلف اندازه گیری می شود. بیشتر مدل هایی که در گروههای بالا جای می گیرند، قادر به رتبه بندی واحدهای کارای غیر راسی نیستند. در کل اندازه گیری کارایی فنی فارل به دو روش محاسبه می شود یکی روش غیر پارامتری و غیر تصادفی است که بر اساس برنامه ریزی ریاضی پایه گذاری شده (CCR) و دیگری روش پارامتری و تصادفی است که روی تکنیک های اقتصادی بنا می شود. هر دو روند نواقص مشترکی دارند و آن اینکه در هر روش معین کردن

نمرات کارایی دقیق آماری مشکل و یا غیر ممکن است. در مدل اقتصادی روش های غیر خطی زیادی برای محاسبه ی نمرات کارایی بکار برده می شود که اکثرا از روی تخمین مدل های پارامتری به دست می آیند. در روند برنامه ریزی ریاضی نیز از آنجایی که روش غیر پارامتری هست، توزیع نمرات کارایی نه شناخته شده است، و نه معین. در صورتی که مشاهده ی اندازه ی دقت آماری برای نمرات کارایی و حدود اطمینان آن ها برای تصمیم گیرنده خیلی مفید و سازنده است. حال نشان می دهیم که چگونه با استفاده از بوت استرپ می توان رتبه بندی نمرات کارایی تولید شده از روند برنامه ریزی خطی را انجام داد.

### ۳ روش بوت استرپ برای رتبه بندی

پشت پرده ی بوت استرپ بر این عقیده استوار است که از روی داده های اصلی، نمونه های جدید کاذب ایجاد می کند، تا توزیع نمونه را تخمین بزند. بوت استرپ به عوض تحلیل تئوری از قدرت محاسباتی استفاده می کند. برای بوت استرپ نمرات کارایی به طور ضمنی از برنامه ریزی خطی استفاده می کنیم. انگیزه ی استفاده ی ضمنی از برنامه ریزی خطی در این است که نمرات کارایی تولید شده از این طریق اندازه هایی نسبی هستند و چون مرز تابع تولید شناخته شده نیست، در عوض کارایی نسبت به بهترین نمونه ی مرزی به دست آمده، از داده های مشاهده شده، سنجیده می شود. و هر گونه تغییر در مشاهدات، متشابهها مجموعه ی نمرات کارایی و در نتیجه مرز بهترین نمونه را تغییر می دهد. پس با استفاده از بوت استرپ می توان علاوه بر ساختن بازه های اطمینان، انحراف نمرات کارایی برنامه ریزی خطی، حدود مرز اقتصادی و یا مرز مجموعه ی امکان تولید را نیز به دست آورد. بوت استرپ یک روند غیر پارامتری برای نتایج آماری می باشد. و از آنجائیکه خودش شبیه روش برنامه ریزی خطی است لذا هیچگونه اعمال نفوذی روی قالب توزیع نمرات کارایی ندارد.

اگر  $\{\rho_1^*, \rho_2^*, \dots, \rho_n^*\}$  مجموعه ی نمرات کارایی به دست آمده از حل مدل BCC فوق باشد، هدف به دست آوردن نمونه های بوت استرپ این نمرات کارایی، اندازه های دقت، چگونگی توزیع و رتبه بندی آن ها می باشد.

در این بحث نمرات کارایی به عنوان اندازه های ورودی اساسی مورد توجه قرار می گیرد و بوت استرپ روی نمرات کارایی اصلی اعمال می شود و فقط ورودیها هستند که در نمونه ی جدید تعدیل می شوند. در نمونه ی جدید، داده ها عبارتند از مقادیر خروجی اصلی برای همه  $n$  تا DMU و ورودیهای تعدیل شده ی تمام DMU ها به غیر از DMU ای که کارایی اش مورد بررسی است. ( در این مورد ورودی اصلی را منظور می کنیم).

نمونه ی جدید از جابجایی  $n-1$  نمره ی کارایی اصلی که از حل مدل BCC به دست آمده اند، حاصل می شود. به طوری که  $t$  امین نمره ی کارایی از مجموعه ی نمرات کارایی اصلی  $P^* = \{\rho_1^*, \dots, \rho_t^*, \dots, \rho_n^*\}$  در جای خودش ثابت باقی می ماند. و مجموعه ی جدید نمرات کارایی  $P^b = \{\rho_1^b, \rho_2^b, \dots, \rho_n^b\}$  که  $\rho_j^b \in P^*$  ( $j = 1, \dots, n$ ) برای ساختن یک تکنولوژی بازگشتی استفاده می شود. برای این منظور روند بوت استرپ را میتوان در چهار مرحله ی زیر اجرا کرد:

۱. حل مدل BCC برای ورودی و خروجی  $X$  و  $Y$  و به دست آوردن نمرات کارایی:

$$P^* = \{\rho_1^*, \dots, \rho_t^*, \dots, \rho_n^*\}$$

۲. تعدیل ماتریس ورودی  $X$  توسط ماتریس  $D$  به  $X^* = D \times X$  (ماتریس  $D$  یک ماتریس قطری است که روی قطر اصلی آن به ترتیب نمرات کارایی اصلی قرار گرفته اند.)

۳. برای هر  $DMU_t$  ( $t = 1, 2, \dots, n$ )

الف. جابجایی  $n-1$  نمره ی کارایی از مجموعه ی  $P^*$  و به دست آوردن نمونه ی جدید از نمرات کارایی به صورت  $P_t(b) = \{\rho_1^b, \rho_2^b, \dots, \rho_t^*, \dots, \rho_n^b\}$  که  $\rho_j^b \in P^*$  یعنی تمام  $\rho_j^b$  ها از نمرات کارایی به دست آمده از مرحله ی اول هستند و به غیر از  $t$  امین نمره ، بقیه جابجا شده اند.

ب. ساختن ماتریس ورودی جدید  $X_t(b) = [D_t(b)]^{-1} \times X^*$  که در آن  $D_t(b)$  یک ماتریس قطری  $n \times n$  میباشد که روی قطر اصلی آن مقادیر نمرات کارایی بوت استرپ شده به ترتیب زیر قرار دارند:

$$\{\rho_1^b, \rho_2^b, \dots, \rho_t^*, \dots, \rho_n^b\}$$

طبیعی است که در نمونه ی جدید بعضی از  $DMU$  ها کارا و برخی دیگر نا کارا خواهد بود. و بردار ورودی  $t$  امین  $DMU$  همان بردار اولیه ی  $x^t$  می باشد، یعنی در مولفه ی  $t$  ام  $X_t(b)$  همان  $x^t$  را خواهیم داشت. زیرا عملیاتی که روی این مولفه انجام می شود به صورت زیر است:

$$(1/\rho_t^*) \cdot (\rho_t^* \cdot x^t) = x^t$$

ج. کارایی تکنیکی  $t$  امین  $DMU$  را نسبت به نمونه ی جدید با متغیرهای  $X_t(b)$  و  $y$  و با استفاده از مدل BCC محاسبه می کنیم. و نمره ی بوت استرپ شده ی  $\rho_t^*(b)$  را به دست می آوریم. حال اگر مراحل الف تا ج را  $B$  بار تکرار کنیم نمرات کارایی بوت استرپ شده ی  $\{\rho_t^*(1), \dots, \rho_t^*(B)\}$  برای  $t$  امین  $DMU$  به دست می آید.

۴. توزیع تجربی نمرات کارایی اصلی  $\rho_t^*$   $t = 1, \dots, n$  می تواند با استفاده از توزیع نمرات کارایی بوت استرپ شده ی  $\rho_t^*(B) = \{\rho_t^*(1), \dots, \rho_t^*(B)\}$   $t = 1, \dots, n$  تخمین زده شود.

اگر فرض کنیم بزرگترین تعداد نمونه های جدید  $B$  باشد و کارایی نسبت به هر نتیجه تکنولوژی بازگشتی محاسبه شده باشد، توزیع تجربی برای نمرات کارایی، از نتایج  $B$  نمره ی کارایی ساخته میشود. در کل تعداد  $B \times n$  تکنولوژی بازگشتی و همین تعداد نمره ی کارایی بوت استرپ شده در این روند داریم. پس از نمایش توزیع تجربی نمرات کارایی بوت استرپ شده ی برنامه ریزی خطی میتوان از آن برای ساختن بازه های اطمینان استفاده کرد. به این ترتیب که با مرتب کردن نمرات کارایی بوت استرپ شده و پیدا کردن صدک های مورد نظر، فاصله ی اطمینان صدکی برای آن نمره ی کارایی را به دست آورد. برای به دست آوردن میانگین و انحراف معیار نمرات کارایی بوت استرپ شده می توان از روش مونت کارلو استفاده کرد. همچنین برای به دست آوردن اریبی نمرات کارایی، میتوان اختلاف بین نمرات کارایی اصلی و میانگین بوت استرپ شده ی آن ها را محاسبه کرد. با افزایش  $B$  میتوان میانگین نمرات کارایی بوت استرپ شده را به صورت مجزا به دست آورد و این نمرات

کارایی غیر یکسان، رتبه‌های واحدها را ممکن می‌سازد.

#### ۴ مثال‌های عددی

حال با اجرای مدل‌های مطرح شده، برای داده‌های خاص، نتایج حاصل را با روش بوت استرپ مقایسه می‌کنیم.

جدول ۱. داده‌های ۹ واحد تصمیم‌گیرنده

| $DMU_S$         | ۱   | ۲   | ۳   | ۴ | ۵   | ۶   | ۷   | ۸ | ۹ |
|-----------------|-----|-----|-----|---|-----|-----|-----|---|---|
| Input( $I_1$ )  | ۱/۵ | ۲   | ۱/۲ | ۳ | ۰/۵ | ۲   | ۳   | ۳ | ۴ |
| Input( $I_2$ )  | ۲   | ۱   | ۱/۵ | ۰ | ۳   | ۳   | ۲/۵ | ۴ | ۱ |
| Output( $o_1$ ) | ۳   | ۱/۵ | ۲/۵ | ۴ | ۳   | ۴/۸ | ۳/۵ | ۴ | ۵ |
| Output( $o_2$ ) | ۱   | ۳   | ۵/۵ | ۱ | ۵/۵ | ۴   | ۳   | ۴ | ۵ |

جدول ۲. نتایج مدل‌های AP، مدل بهبود خروجی، روش بوت استرپ

| $DMU_S$ | ۱    | ۲    | ۳      | ۴    | ۵      | ۶      | ۷    | ۸    | ۹      |
|---------|------|------|--------|------|--------|--------|------|------|--------|
| AP      | ۰/۸۴ | ۰/۷۲ | ۱/۵۶   | ۱/۳۱ | ۲/۷۸   | ۰/۹۶   | ۰/۵۹ | ۰/۵۶ | ۱/۰۶   |
| ranking | -    | -    | ۲      | ۳    | ۱      | -      | -    | -    | ۴      |
| IM*     | ۱/۶۷ | ۱/۸۳ | ۱/۰۰   | ۱/۲۵ | ۰/۹۸   | ۱/۰۴   | ۱/۴۳ | ۱/۲۵ | ۰/۸۸   |
| ranking | -    | -    | ۳      | -    | ۲      | -      | -    | -    | ۱      |
| SA**    | ۰/۸۲ | ۰/۷۱ | ۱/۵۴   | ۰/۷۸ | ۲/۶۸   | ۰/۹۴   | ۰/۵۹ | ۰/۵۶ | ۱/۰۴   |
| ranking | -    | -    | ۲      | ۴    | ۱      | ۵      | -    | -    | ۳      |
| BS***   | -    | -    | ۰/۸۵۹۴ | ۱    | ۰/۸۷۶۲ | ۰/۸۵۷۱ | -    | -    | ۰/۸۶۱۵ |
| ranking | -    | -    | ۴      | ۱    | ۲      | ۵      | -    | -    | ۳      |

جدول ۳. داده‌های ۷ واحد تصمیم‌گیری

| $DMU_S$         | $A_1$ | $A_2$ | $A_3$ | B | C  | D  | E  |
|-----------------|-------|-------|-------|---|----|----|----|
| Input( $I_1$ )  | ۲     | ۰/۱   | ۰     | ۵ | ۱۰ | ۱۰ | ۲  |
| Input( $I_2$ )  | ۸     | ۸     | ۸     | ۵ | ۴  | ۶  | ۱۲ |
| Output( $o_1$ ) | ۱     | ۱     | ۱     | ۱ | ۲  | ۲  | ۱  |
| Output( $o_2$ ) | ۲     | ۲     | ۲     | ۱ | ۱  | ۱  | ۲  |

جدول ۴. نتایج حاصل از مدل‌ها بر روی داده‌ها

| $DMU_S$ | CCR | SuperSBM(I) | MAJ  | MMAJ* | JHF   | LJK  | BS     |
|---------|-----|-------------|------|-------|-------|------|--------|
| $A_1$   | ۱   | ۱/۲۵        | ۱/۲۸ | ۱/۱۵  | ۰/۰۴۱ | ۱/۱۷ | ۰/۸۳۱۶ |
| $A_2$   | ۱   | Indefinable | ۱/۳۱ | ۱/۱۷  | ۰/۰۶۲ | ۱/۲۷ | ۰/۸۹۵۵ |
| $A_3$   | ۱   | ۱۰/۷۵       | ۱/۳۱ | ۱/۱۷  | ۰/۰۶۱ | ۱/۲۶ | ۱      |

همان‌طور که از جدول بالا مشاهده می‌شود، مدل رتبه‌بندی SuperSBM(I) در ارزیابی واحد  $A_2$  ناتوان است، همچنین مدل‌های MAJ و MAJ اصلاح شده رتبه‌ی یکسانی را برای  $A_2$  و  $A_3$  به دست می‌دهند. با اینکه مدل JHF واحدهای فوق را به صورت کامل رتبه‌بندی می‌کند، ولی با توجه به نوع وزن‌ها، ممکن است برای داده‌های بزرگ منجر به خطاهای عددی شود. مدل LJK با اینکه رتبه‌بندی را انجام می‌دهد ولی مقادیر

مربوط به  $A_1$  و  $A_2$  خیلی نزدیک به هم هستند و این پایداری مدل را زیر سوال می برد. برای داده های مذکور مدل BS بهترین عملکرد را از خود نشان می دهد، این مدل با به دست دادن نمرات کارایی متمایز، واحدهای تصمیم گیرنده را به طور کامل رتبه بندی می کند.

مثالی برای مقایسه ی مدل های نرم یک، نرم بینهایت، مدل غیر شعاعی MAJ، مدل شعاعی AP و روش بوت استرپ (BS) با داده های زیر را در نظر می گیریم:

جدول ۵. داده های ۶ واحد تصمیم گیری

| $DMU_S$ | ۱  | ۲ | ۳ | ۴  | ۵ | ۶ |
|---------|----|---|---|----|---|---|
| $I_1$   | ۱۳ | ۶ | ۲ | ۱  | ۹ | ۴ |
| $I_2$   | ۱  | ۳ | ۶ | ۱۰ | ۵ | ۸ |
| O       | ۱  | ۱ | ۱ | ۱  | ۱ | ۱ |

همان طور که مشاهده می شود، واحدهای تصمیم گیرنده ی جدول بالا با دو ورودی مجزا، خروجی یکسان یک را تولید می کنند.

جدول ۶. نتایج رتبه بندی

| $DMU_S$    | ۱    | ۲      | ۳      | ۴      | ۵           | ۶           |
|------------|------|--------|--------|--------|-------------|-------------|
| $L_1$      | ۰/۲۰ | ۰/۱۲   | ۰/۱۴   | ۰/۰۷   | inefficient | inefficient |
| ranking    | ۱    | ۳      | ۲      | ۴      | -           | -           |
| $L_\infty$ | ۰/۱۵ | ۰/۰۵   | ۰/۰۷   | ۰/۰۷   | inefficient | inefficient |
| ranking    | ۱    | ۴      | ۳      | ۲      | -           | -           |
| MAJ        | ۱/۲۰ | ۱/۰۷   | ۱/۰۹   | ۱/۰۸   | inefficient | inefficient |
| ranking    | ۱    | ۴      | ۲      | ۳      | -           | -           |
| AP         | ۲/۹۹ | ۱/۲۱   | ۱/۲۹   | ۱/۹۹   | inefficient | inefficient |
| ranking    | ۱    | ۴      | ۳      | ۲      | -           | -           |
| BS         | ۱    | ۰/۹۵۳۰ | ۰/۹۶۷۶ | ۰/۹۵۳۲ | inefficient | inefficient |
| ranking    | ۱    | ۴      | ۳      | ۲      | -           | -           |

جالب اینکه رتبه بندی روش بوت استرپ و AP دقیقاً یکسان هستند.

جدول ۷. داده های ۷ واحد تصمیم گیری

| $DMU_S$         | ۱    | ۲    | ۳    | ۴    | ۵    | ۶    | ۷   | ۸    | ۹    | ۱۰   | ۱۱   |
|-----------------|------|------|------|------|------|------|-----|------|------|------|------|
| Input( $I_1$ )  | ۱۶۲۱ | ۲۷۱۸ | ۱۵۲۳ | ۵۵۱۴ | ۱۹۴۱ | ۱۴۹۶ | ۹۳۲ | ۲۰۱۳ | ۱۸۹۱ | ۲۲۷۷ | ۱۹۹۵ |
| Input( $I_2$ )  | ۴۳۶  | ۳۱۴  | ۳۴۵  | ۱۳۱۴ | ۵۰۷  | ۳۲۱  | ۱۵۸ | ۱۰۳۷ | ۹۷۶  | ۸۹۱  | ۶۹۳  |
| Input( $I_3$ )  | ۲۰۵  | ۲۲۱  | ۲۱۵  | ۵۵۳  | ۳۰۹  | ۳۳۹  | ۲۰۰ | ۴۱۲  | ۳۹۹  | ۴۱۸  | ۳۴۹  |
| Output( $O_1$ ) | ۱۷۴  | ۱۷۲  | ۱۶۰  | ۴۸۷  | ۲۲۰  | ۱۰۹  | ۳۷  | ۱۹۸  | ۱۹۱  | ۲۴۱  | ۱۶۷  |
| Output( $O_2$ ) | ۴۹۷  | ۴۹۷  | ۴۴۳  | ۱۹۲۵ | ۵۲۱  | ۶۹۹  | ۴۳۱ | ۴۷۱  | ۴۹۱  | ۳۷۹  | ۴۱۲  |
| Output( $O_3$ ) | ۲۲   | ۲۲   | ۲۲   | ۶۳   | ۳۶   | ۳۸   | ۱۹  | ۳۲   | ۲۲   | ۲۸   | ۳۱   |

روش های رتبه بندی با مجموعه وزن های مشترک را با روش بوت استرپ مقایسه می کنیم:

جدول ۸. روش رتبه‌بندی با مجموعه وزن‌های مشترک با روش بوت استرپ

| $DMU_S$                 | ۱      | ۲      | ۳      | ۴    | ۵      | ۶      | ۷      | ۸    | ۹    | ۱۰   | ۱۱   |
|-------------------------|--------|--------|--------|------|--------|--------|--------|------|------|------|------|
| $Efficiency (\theta^*)$ | ۱      | ۱      | ۱      | ۱    | ۱      | ۱      | ۱      | ۰/۸۷ | ۰/۹۱ | ۰/۹۳ | ۰/۷۹ |
| CSW <sub>1</sub>        | ۰/۹۷   | ۰/۸۴   | ۰/۹۸   | ۱    | ۱      | ۱      | ۰/۸۶   | ۰/۶۲ | ۰/۵۷ | ۰/۶۴ | ۰/۶۹ |
| ranking                 | ۳      | ۵      | ۲      | ۱    | ۱      | ۱      | ۴      | ۸    | ۹    | ۷    | ۶    |
| CSW <sub>2</sub>        | ۱/۲۹   | ۰/۷۱   | ۱/۳۵   | ۰/۸۷ | ۱/۲۹   | ۰/۸۵   | ۰/۷۱   | ۰/۷۱ | ۰/۷۱ | ۰/۷۷ | ۰/۷۲ |
| ranking                 | ۲      | ۷      | ۱      | ۳    | ۲      | ۴      | ۷      | ۷    | ۷    | ۵    | ۶    |
| CSWA                    | ۱      | ۰/۶۹   | ۰/۹۹   | ۰/۹۷ | ۱      | ۱      | ۰/۸۲   | ۰/۷۲ | ۰/۷۳ | ۰/۷۴ | ۰/۷۱ |
| ranking                 | ۱      | ۱۱     | ۴      | ۵    | ۳      | ۲      | ۶      | ۷    | ۸    | ۹    | ۱۰   |
| AP                      | ۱/۰۸   | ۱/۲۰   | ۱/۰۷   | ۱/۴۴ | ۱/۱۴   | ۱/۲۷   | ۱/۲۵   | ۰/۸۷ | ۰/۹۱ | ۰/۹۳ | ۰/۷۹ |
| ranking                 | ۶      | ۴      | ۷      | ۱    | ۵      | ۲      | ۳      | —    | —    | —    | —    |
| BS                      | ۰/۹۷۶۹ | ۰/۹۹۷۰ | ۰/۹۷۲۱ | ۱    | ۰/۹۸۳۲ | ۰/۹۹۹۰ | ۰/۹۹۸۴ | —    | —    | —    | —    |
| ranking                 | ۶      | ۴      | ۷      | ۱    | ۵      | ۲      | ۳      | —    | —    | —    | —    |

همان‌طور که از نتایج مدل CWA برمی‌آید، بعضی از واحدهای ناکارا دارای رتبه‌ی بالاتری نسبت به واحدهای کارا هستند، اینگونه واحدها، هر چند ممکن است در مولفه‌هایی از ورودی (خروجی) دارای کمترین (بیشترین) مقدار باشند، ولی در عمل، در سطح واحدهای ناکارا و یا حتی کمتر از آن‌ها بازدهی دارند. در صورتی که چنین مشکلی در روش بوت استرپ وجود ندارد و نتایج رتبه‌بندی حاصل از روش بوت استرپ و AP دقیقاً یکسان هستند.

با استفاده از داده‌های زیر روش بوت استرپ (BS)، را با مدل‌های AP، MAJ، و مدل مونت کارلو مقایسه می‌کنیم:

جدول ۹. داده‌ها

| $DMU_S$ | A | B | C | D | E | F |
|---------|---|---|---|---|---|---|
| Input   | ۱ | ۳ | ۴ | ۹ | ۳ | ۶ |
| Output  | ۱ | ۳ | ۴ | ۷ | ۲ | ۵ |

جدول ۱۰. رتبه‌بندی بر روی داده‌های جدول ۹

| $DMU_S$                    | A      | B      | C       | D           | E           | F           |
|----------------------------|--------|--------|---------|-------------|-------------|-------------|
| AP                         | ۲/۲۹   | ۱/۰۰   | ۱/۱۳    | ۱/۰۰        | ۰/۶۷        | ۰/۹۴        |
| ranking                    | ۱      | ۳      | ۲       | ۳           | —           | —           |
| MAJ                        | ۱/۰۰   | ۱/۰۰   | ۱/۰۰    | ۰/۷۸        | ۰/۸۹        | ۰/۸۹        |
| ranking                    | ۱      | ۱      | ۱       | —           | —           | —           |
| Montcarlo(N <sub>H</sub> ) | ۱۴۲۰   | ۷۱۶۸   | ۱۴۲۰۱   | ۲۸۵۰۴       | Inefficient | Inefficient |
| ranking                    | ۴      | ۳      | ۲       | ۱           | —           | —           |
| BS                         | ۰/۷۹۳۰ | ۰/۷۷۹۱ | ۰/۷۸۷۶۱ | Inefficient | Inefficient | Inefficient |
| ranking                    | ۱      | ۳      | ۲       | —           | —           | —           |

روشن است که در این مثال روش MAJ قادر به رتبه‌بندی نیست. و نتایج حاصل از روش بوت استرپ و AP

دقیقا یکسان هستند.

## ۵ نتیجه گیری

روش های رتبه بندی مطرح شده در این مقاله معایب و محاسنی داشتند که به طور خلاصه در ذیل می آوریم. مدل AP کمکی برای رتبه بندی واحدهای کارای غیر راسی نمی نماید. در ماهیت ورودی برای DMUهای با ورودی های خاص، ممکن است نشدنی گردد و همچنین در بعضی از موارد ممکن است، تغییرات کوچکی در داده ها، تغییرات زیادی را در  $\theta^*$  حاصل کند. در این مدل حرکت به سوی مرز در امتداد شعاعی صورت می گرفت که ممکن بود سطح پوششی PPS را قطع نکند. در این صورت مساله جواب شدنی ندارد و یا ممکن است در فاصله ی بسیار دور، سطح پوششی PPS را قطع نماید. در این صورت مساله ناپایدار خواهد بود. در مدل MAJ حرکت به سوی مرز در امتداد محورهای ورودی -در ماهیت ورودی- و با قدم های مساوی انجام می شود. ممکن است در بعضی حالات مدل MAJ نشدنی شود. و این مدل نیز برای واحدهای کارای غیر راسی پیشنهادی نمی دهد.

مدل SBM هم برای واحدهای کارای غیر راسی پیشنهادی ندارد. از مهمترین خصوصیت این مدل آن است که اندازه ی ابر کارایی  $\delta^*$  نسبت به تغییر واحد پایدار است.

مدل LJK هم در ارزیابی عملکرد واحدهای تصمیم گیرنده در مقابل تغییر واحد پایدار است. مانند مدل های نرم یک نرم بینهایت و بقیه مدل ها برای رتبه بندی واحدهای کارای غیر راسی پیشنهادی ندارد. در صورتی که روش بوت استرپ مطرح شده در این مقاله می تواند رتبه بندی واحدهای کارای راسی و غیر راسی را انجام دهد. روند بوت استرپ به عوض تحلیل تئوری از قدرت محاسباتی استفاده می کند. برای بوت استرپ نمرات کارایی به طور ضمنی از برنامه ریزی خطی استفاده می کند. و به همین دلیل هیچکدام از مشکلات روش های قبلی را ندارد. تنها ایرادی که می توان به این روش گرفت آن است که حجم محاسبات خیلی زیاد است، که این نیز به کمک کامپیوترهای پر سرعت قابل رفع می باشد.

## منابع

- [1] Anderson P., Peterson N. C., (1993), *A procedure for ranking efficient units in data envelopment analysis*. Management science. 39 (10), 1261-1264.
- [2] Banker R. D., Charnse A., Cooper, W. W., (1984). *Models for the estimation of the technical and scale inefficiencies in data envelopment analysis*. Management science. 30, 1078-1092.
- [3] Charnes, A., Cooper, W. W., Rhodes E., (1978). *Measuring the efficiency of desicio making units*. European Journal of Operation Research. 2, 429-444.
- [4] Cooper, W. W., Park, K. S., Pastor J.T., (1999). *RAM: a range adjusted measure of inefficiency for use with additive woodless, and relations to other models and measures in DEA*. Journal of Productivity Analysis. 11, 5-24.
- [5] Hosseinzadeh Lotfi, F., Jahanshahloo, G.R., Memariani, A., (2000). *A method for finding common set of weights by multiple objective programming in data envelopment analysis*. South West jurnal of pure and applied mathematics. 1, 44-54.
- [6] Jahanshahloo, G.R., Afzalinejad, M., (2006). *A ranking method based on a full inefficient frontier*". Applied Mathematical modeling. 30, 248-260.
- [7] Jahnsahlloo, G.R., Hosseinzadeh Lotfi, F., Rezaie, V., Khanmohamadi M., (2011). *Ranking DMUs by ideal points with interval data in DEA*. Applied Mathematical Modeling. 35, 218-229.

- [8] Jahanshahloo, G.R., Hosseinzadeh Lotfi, F., Rezai Balf F., Zhiani Rezai H., (2008). *Using mont carlo method for ranking efficient units* ". Applied Mathematics and Computation. 201, 613-620.
- [9] Jahanshahloo G. R., Hosseinzadeh Lotfi, F., Rezai Balf F., Zhiani Rezai, H., Akbarian D. (2004). *Ranking efficient DMUs using tchebycheff norm*. Working Paper.
- [10] Jahanshahloo, G. R., Hosseinzadeh Lotfi, F., Sanaei, M., Fallah Jelodar, M., (2008). *Review of ranking models in the data envelopment analysis* ". Applied Mathematical Sciences, Vol. 2, no. 29, 1431-1448.
- [11] Jahanshahloo, G. R., Memariani, A., Hosseinzadeh Lotfi, F., Rezai, H. Z., (2005). *A note on som DEA models and finding efficiency and complete ranking using common set of weights*. Applied Mathematics and Compiotation. 166, 265-281.
- [12] Jahanshahloo, G. R., Hosseinzadeh Lotfi, F., Shoja, N., Tohidi, G., Razavian, S. (2004). *Ranking using l1-norm in data envelopment analysis*. Applied Mathematics and Compotation. 153, 215-224.
- [13] Jie W., Feng Y., Liang L., (2010). *A modified complete ranking of DMUs using restrctions in DEA models*. Applied Mathematcsand ComputationI. 217, 745-751.
- [14] Li S., Jahanshahloo, G. R., Khodabakhshi, M. (2007). *A super-efficiency model for ranking units in data envelopment analysis*. Applied Mathematics and Computation. 184, 638-648.
- [15] Liu, F-H-F., Peng, H -H. (2006). *Ranking of units on the DEA frontier with common weights*. Computers and operations Research. 35 (5), 1624-1637.
- [16] Mehrabian, S., Alirezaee, M. R., Jahanshahloo, G. R., (2002), " *A complete efficiency ranking of decision making units in data envelopment analysis*. Computational optimization and application 4, 261-266.
- [17] Saati, M. S., Zarafat Angiz, M., Jahanshahloo, G. R., (1999). *A model for ranking decision making units in data envelopment analysis*. Recerca Oprerativa vol.31, no 97, 47-59.
- [18] Sueyoshy, T., (1999). *DEA nonparametric ranking test and index measurment: Slack Adjusted DEA and an application to Jpanese agriculture cooperative*. Omega, INT .J. MGMT. SCI 27, 315-326.