

استفاده از رهیافت تحلیل پوششی داده‌های بوت‌استرپ و الگوریتم LSW در راستای سنجش ضریب کارایی صنایع کارخانه‌ای ایران

محمدنبی شهیکی تاش^{۱*}، جواد شهرکی^۱، مصطفی خواجه‌حسینی^۲

۱- دانشیار، دانشگاه سیستان و بلوچستان، گروه اقتصاد، زاهدان، ایران

۲- دانشجوی دکتری، دانشگاه سیستان و بلوچستان، گروه اقتصاد، زاهدان، ایران

رسید مقاله: ۴ تیر ۱۳۹۶

پذیرش مقاله: ۲۰ آذر ۱۳۹۶

چکیده

این مطالعه با استفاده از رویکرد تحلیل پوششی داده‌های بوت‌استرپ نهاده محور با استفاده از الگوریتم LSW، به تخمین مقادیر کارایی تورش اصلاح‌شده کارایی فنی، مدیریتی و مقیاس کارگاه‌های صنعتی ۱۰ نفر کارکن و بیش‌تر در سال ۱۳۹۰ پرداخته است. بوت‌استرپ یک تکنیک بازنمونه‌گیری است که برای تخمین خواص توزیع نمونه‌گیری یک تخمین‌زننده به کار گرفته می‌شود. دلیل استفاده از این روش در این مقاله وابستگی و حساسیت مقادیر کارایی و رتبه‌بندی هر بنگاه اقتصادی به تغییرات ترکیب نمونه و لزوم اصلاح تورش ایجادشده ناشی از محاسبه مقادیر کارایی در روش‌های معمول تحلیل پوششی داده‌ها است. نتایج به‌دست‌آمده از این مطالعه حاکی از آن است که میانگین کارایی مدیریتی و مقیاس بخش صنعت کشور به ترتیب ۷۷ و ۹۴ درصد است؛ این امر نشان‌دهنده آن است که مهم‌ترین علت ناکارایی فنی بخش صنعت ناکارآمدی مدیریتی می‌باشد. به این ترتیب تنها ۹ درصد از صنایع کشور به لحاظ فنی به صورت کاملاً کارآمد عمل نموده و ۴۸ درصد از آن‌ها کارایی فنی پایین‌تر از ۷۰ درصد دارند.

کلمات کلیدی: تحلیل پوششی داده‌ها، بازنمونه‌گیری بوت‌استرپ، بخش صنعت، کارایی، منابع طبیعی، اصلاح تورش.

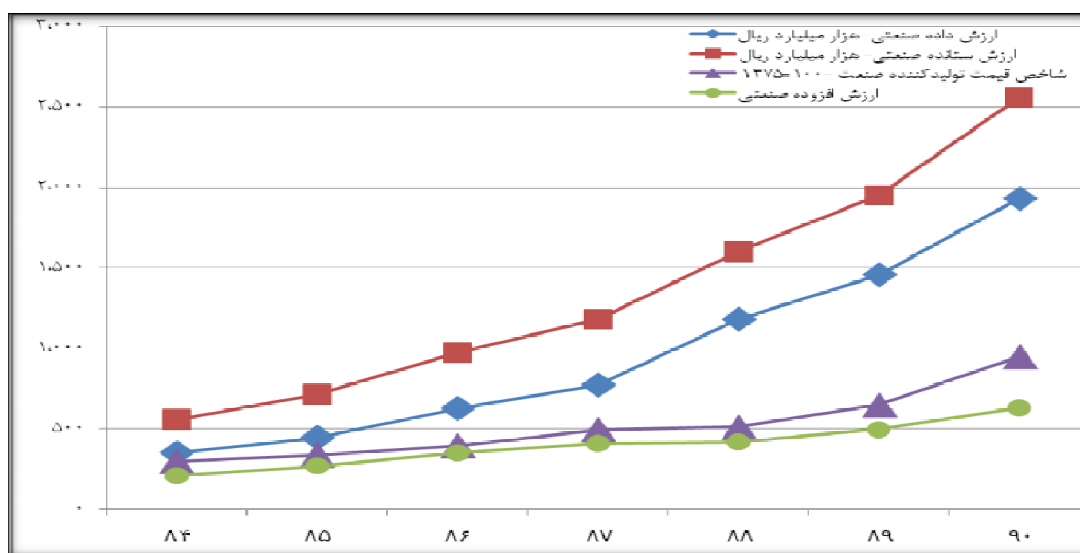
۱ مقدمه

بخش صنعت کشور با ۱۴۹۶۲ کارگاه ۱۰ نفر کارکن و بیش‌تر با به کارگیری ۱۲۴۲۹۸۳ نفر نیروی کار و ارزش افزوده‌ای معادل ۶۲۹۲۴۹ میلیارد ریال در سال ۱۳۹۰، یکی از مهم‌ترین بخش‌های اقتصادی کشور محسوب شده و سهم بالایی از تولید ناخالص داخلی را به خود اختصاص می‌دهد. همان‌طور که شکل زیر نشان می‌دهد ارزش ستانده‌های بخش صنعتی کشور علیرغم عدم استفاده از تکنولوژی‌های جدید در طول سالیان ۹۰-۱۳۸۴

* عهده‌دار مکاتبات

آدرس الکترونیکی: mohammad_tash@yahoo.com

دارای رشدی صعودی بوده و این نشان‌دهنده اهمیت و جایگاه ویژه و استراتژیک این بخش برای برنامه‌ریزان و مدیران کشور است [۱].



شکل ۱. روند ارزش داده و ستانده بخش صنعت کشور در طی سال‌های ۹۰-۱۳۸۴

در عصر حاضر، به دلیل جهانی‌شدن اقتصادها، بخش صنایع کشورها در حال رقابت شدید با یکدیگر برای تصاحب بازارهای جهانی هستند و همواره تلاش می‌کنند که با بهبود مدیریت، تکنولوژی و تولید در مقیاس مناسب وضعیت صنایع خود را در این رقابت تنگاتنگ جهانی ارتقا دهند. از آنجاکه یکی از ملزومات رشد و توسعه صنایع مذکور، دسترسی مناسب و کافی به نهاده‌های تولید مانند انرژی و مواد اولیه است و استفاده بیش‌تر از آنها با افزایش آلودگی محیط زیستی و هزینه‌های نهایی فزاینده ناشی از استخراج و واردات آنها همراه است، ضرورت دارد که در راستای صرفه‌جویی و استفاده کارآمد از این نهاده‌های ارزشمند، همواره روند شاخص‌های مختلف ارزیابی کارایی و بهره‌وری را به‌طور مداوم پیگیری کرد. با توجه به آنچه ذکر شد، یکی از الزامات مهم و اساسی ارتقای توان رقابت صنایع و بنگاه‌های زیرمجموعه آنها، بهبود رتبه کارایی و در نتیجه افزایش بهره‌وری در استفاده از عوامل تولید است. پرداختن به مساله کارایی برای همه کشورها الزامی بوده و این نیاز برای کشورهای در حال توسعه بیش‌تر است؛ زیرا این کشورها به علت نداشتن تکنولوژی برتر به اتلاف بیش‌تر منابع و نهاده‌های تولید می‌پردازند و امکان اینکه کارایی در این کشورها خیلی پایین باشد، بیش‌تر است [۲].

از طرف دیگر در یک نگاه کلی، اقتصاد کشور ما به‌شدت به درآمدهای ارزی حاصل از صادرات نفت خام و واردات نهاده‌های تولید و مواد اولیه وابسته است و باید برای عملیاتی کردن نظریه اقتصاد مقاومتی و کاهش وابستگی به خارج تلاش نماییم؛ بنابراین بهترین و عمده‌ترین راهکار برای تولید کالاها و خدمات، ارتقای سطح کارایی و بهره‌وری است که باید به‌طور جدی به آن پرداخته شود. با ارتقای شاخص‌های یادشده می‌توان

محصولات صنعتی با کیفیت برتر و قیمتی نازل تر تولید و با ارتقای سطح رقابت پذیری زمینه های لازم جهت ورود به بازارهای جهانی را فراهم نمود.

با توجه به اهمیت موضوع همان طور که در ادامه به آن پرداخته می شود، این مقاله کارایی صنایع کارخانه ای ایران را مورد بررسی قرار داده است. در این مطالعه پس از ارائه مقدمه ابتدا در قسمت دوم به مهم ترین مطالعات مرتبط داخلی و خارجی پرداخته شده و پس از آن در بخش سوم مبانی نظری و روش تحقیق به کار گرفته شده شرح داده می شود. پس از ارائه اطلاعات لازم درباره داده های مورد استفاده در این مطالعه در بخش چهارم، بخش پنجم به بیان نتایج به دست آمده از تخمین مدل مورد نظر می پردازد. در نهایت این مقاله با ارائه بحث و نتیجه گیری و همچنین ارائه پیشنهادها در بخش ششم پایان می یابد.

۲ ادبیات تحقیق

با توجه به اهمیت ارزیابی کارایی صنایع، تا به حال با به کارگیری روش های مختلف مطالعات بی شماری راجع به این موضوع انجام شده است. در این قسمت ابتدا به مهم ترین مطالعات تجربی خارجی و سپس به مهم ترین مطالعات داخلی پرداخته می شود.

هاکس و تزریمس [۳] در مطالعه خود با استفاده از رهیافت تحلیل پوششی داده های بوت استرپ^۱ و به کارگیری داده های مالی به ارزیابی کارایی ۲۳ بخش تولیدی پرداختند؛ در مرحله اول، نتایج این مطالعه نشان می دهد که تحلیل حساسیت نمرات کارایی به دست آمده تورش دار هستند سپس در مرحله دوم با به کارگیری تکنیک بوت استرپ نتایج به دست آمده به طور قابل توجهی بهبود یافتند.

لی و ورسینکتن [۴] در مطالعه خود به وسیله داده های نمونه گیری شده در سال های ۲۰۰۱ تا ۲۰۰۵ و با استفاده از تکنیک تحلیل پوششی بوت استرپ (دومرحله ای)^۲ به اندازه گیری کارایی فنی خطوط هوایی داخلی و بین المللی آمریکا و سایر کشورها پرداختند. نتایج این مطالعه گویای آن است عملکرد خطوط هواپیمایی بین المللی غیر اروپایی و آمریکایی به ویژه هواپیمایی سنگاپور و به میزان کمتری جی ای ال کاتای پاسیفیک در سطحی از کارایی است که می توان از آن به عنوان معیاری برای سنجش بهبود عملکرد و کارکرد خطوط هواپیمایی آمریکایی و اروپایی که عملکردی ضعیف دارند، استفاده کرد.

ذکریا، صالح و همکاران [۵] با استفاده از تحلیل پوششی داده های بوت استرپ به بررسی کارایی ۱۲ بانک اسلامی در مالزی پرداختند. نتایج این تحقیق نشان می دهد که به صورت کلی استفاده از این روش نسبت به روش سنتی که دارای تورش می باشد، دقیق تر است.

مورنو و همکاران [۶] با تخمین تابع تولید کاب-داگلاس مرزی تصادفی به ارزیابی کارایی فنی صنعت منسوجات اسپانیایی در سال های ۲۰۰۹-۲۰۰۲ پرداختند. نتایج این مطالعه نشان می دهد که بنگاه های بررسی شده

^۱ Bootstrapping Data Envelopment Analysis (BDEA).

^۲ DEA double bootstrapping model.

طی سال‌های ۲۰۰۲ تا ۲۰۰۵ از سطوح کارایی فنی بالاتری نسبت به سال‌های ۲۰۰۵ تا ۲۰۰۹ برخوردار بوده‌اند و این نتایج روند رو به زوالی برای بنگاه‌های مورد بررسی نشان می‌دهد.

علاوه بر مطالعات انجام شده در بخش صنعت که در بالا به آن‌ها پرداخته شد، می‌توان به مطالعاتی که با روشی مشابه این تحقیق انجام شده‌اند نیز مانند مطالعات برومر [۷]، گوچت و بالکومب [۸]، دونگ [۹]، بالکومب و همکاران [۱۰] و ادک [۱۱] در سایر بخش‌ها اشاره نمود. از مهم‌ترین مطالعات داخلی برای ارزیابی کارایی در بخش صنعت نیز می‌توان به مطالعات زیر اشاره کرد.

زاراء نژاد و همکاران [۱۲] در مطالعه خود با استفاده از مدل اثرات ناکارایی بتیس و کوئلی میزان کارایی فنی و صنایع کارخانه‌ای ایران را طی سال‌های ۸۶-۱۳۷۵ مورد ارزیابی قرار دادند. نتایج تحقیق آن‌ها نشان می‌دهد که میانگین کارایی فنی صنایع مورد بررسی ۵۵ درصد است. همچنین صنایع فعال در زمینه تولید محصولات اساسی مسی و تولید فراورده‌های نفتی تصفیه شده به ترتیب با سطح کارایی ۸۳ و ۷۸ درصد به طور نسبی از سطح کارایی فنی بالاتری در مقایسه با دیگر فعالیت‌های صنعتی برخوردار هستند. در مقابل صنایع فعال در زمینه تولید آجر، آماده‌سازی و آرد کردن غلات و حبوبات به ترتیب با سطوح کارایی ۲۱ و ۲۳ درصد پایین‌ترین میزان کارایی فنی را به خود اختصاص داده‌اند.

آزادی نژاد و همکاران [۲] در مطالعه خود با استفاده از رهیافت تحلیل پوششی داده‌ها به بررسی عوامل مؤثر در کارایی فنی بخش صنعت استان‌های کشور پرداختند. آن‌ها ابتدا هر کدام از ۲۸ استان کشور را به منزله یک واحد تصمیم‌گیر در نظر گرفتند و سپس به اندازه‌گیری کارایی فنی استان‌های کشور در سال‌های ۱۳۷۵ تا ۱۳۸۶ پرداختند. نتایج مطالعه آن‌ها نشان می‌دهد که استان‌های مرکزی، بوشهر، کرمان، هرمزگان، تهران و خوزستان در بخش صنعت نسبت به سایر استان‌ها از کارایی فنی بالاتری برخوردارند.

شهیکی تاش و همکاران [۱۳] در مطالعه خود با رهیافت تابع مرزی تصادفی به ارزیابی ناکارایی فنی صنایع کارخانه‌ای ایران در طی سال‌های ۱۳۷۴ تا ۱۳۸۸ پرداختند. نتایج کار آن‌ها نشان می‌دهد که بیش‌ترین ضریب ناکارایی فنی به ترتیب مربوط به صنایع تولید پوشاک و عمل آوردن و رنگ کردن پوست خردار، انتشار و چاپ و تکثیر رسانه‌های ضبط شده و ساخت منسوجات و بازیافت می‌باشد. در مقابل بیش‌ترین کارایی فنی نیز مربوط به صنعت ساخت مواد و محصولات شیمیایی است.

۳ مبانی نظری و روش تحقیق

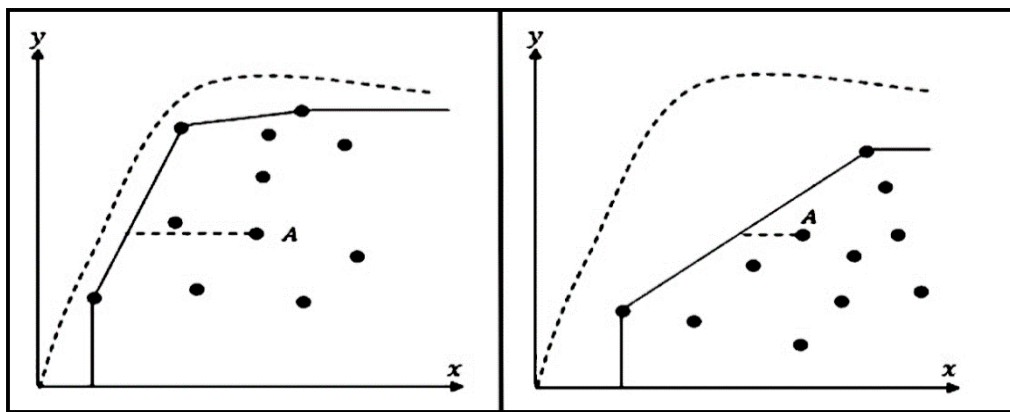
تا قبل از سال ۱۹۷۸ تحقیقات فراوانی برای محاسبه کارایی واحدهای تصمیم‌گیرنده یک سیستم صورت گرفته بود. عمده این تحقیقات منجر به ایجاد روش‌های پارامتریک گردیدند. این روش‌ها در برخی حالات خاص کارساز بودند؛ ولی در حالت کلی دو مشکل عمده نظری و کاربردی، استفاده از آن‌ها در سطح گسترده را غیرممکن می‌ساخت. این دو مشکل عبارت‌اند از:

الف- روش‌های پارامتریک برای حالت‌های یک یا چند ورودی و تنها یک خروجی مناسب هستند.

ب- محاسبه پارامترها و تعیین تابع پارامتریک در حالت کلی آسان نیست و نیاز به تصریح فرم تبعی خاص دارد.

در سال ۱۹۵۷ فارل [۱۴] طی مقاله‌ای روش اندازه‌گیری کارایی را بر مبنای تئوری‌های اقتصادی معرفی و کارایی بخش کشاورزی آمریکا را به روش غیر پارامتریک محاسبه نمود. فارل با استناد بر اصول پنج‌گانه، مجموعه‌ای به نام «مجموعه امکانات تولید» ساخت و قسمتی از مرز آن را به عنوان تخمینی از تابع تولید در نظر گرفت. در روش فارل هر واحد تصمیم‌گیرنده‌ای که روی این مرز قرار گیرد کارا می‌باشد و در غیر این صورت ناکارا تلقی می‌گردد.

به دلیل مشکلات علمی در اندازه‌گیری و محدودیت‌هایی که در روش فارل مطرح بود، این روش کاربرد عملی چندانی نیافت و سال‌ها مسکوت ماند، تا اینکه در سال ۱۹۷۸ چارنز، کوپر، و رودس [۱۵] (CCR) با جامعیت بخشیدن به روش فارل، به گونه‌ای که خصوصیت فرآیند تولید با چند عامل تولید و محصول را در برگیرد، روش تحلیل پوششی داده‌ها^۱ (DEA) را معرفی نمودند. در این روش برای تخمین تابع تولید به پیش‌فرض خاصی در مورد شکل تابع نیاز نبوده و کارایی یک بنگاه نسبت به کارایی سایر بنگاه‌ها اندازه‌گیری می‌شود. برای آشنایی با این روش، فرض کنید سیستم تحت ارزیابی شامل n واحد تصمیم‌گیرنده^۲ $(DMU_1, DMU_2, \dots, DMU_n)$ باشد که هر DMU_j ، m ورودی $X = (x_{1j}, x_{2j}, \dots, x_{mj})$ را برای تولید s خروجی $Y = (y_{1j}, y_{2j}, \dots, y_{sj})$ مصرف می‌نماید [۱۴ و ۱۵]. متأسفانه در دنیای واقعی به علت استفاده از نمونه، قادر به محاسبه بیش‌ترین میزان تولید شده از یک ورودی مشخص نیستیم؛ زیرا ما تنها یک نمونه از یک جامعه ناشناخته را در اختیار داریم و در این صورت مرز کارایی تولید جامعه، نامشخص خواهد بود. همان‌طور که در شکل‌های زیر نشان داده شده است، کارایی واحد A در شکل (۲) نسبت به شکل (۳) به‌طور محسوسی تغییر یافته است.



شکل ۳. نمونه دوم [۱۶].

شکل ۲. نمونه اول [۱۶].

این تغییر در مقدار کارایی ناشی از طبیعت ناپارامتریک مدل DEA است [۱۶]. همان‌طور که مشاهده می‌شود، مرز DEA وابسته به نمونه بوده و به آن حساس می‌باشد، به‌طوری‌که با تغییر نمونه مرز قبلی فرو می‌پاشد؛ البته تمام

^۱ Data Envelopment Analysis.

^۲ Decision Making Units.

ضعف این مدل به دلیل ناپارامتری بودن آن نیست؛ بلکه به‌اندازه نمونه نیز بستگی دارد. برای حل مشکل یاد شده، سیمار [۱۷]، روشی تحت عنوان تحلیل پوششی داده‌های بوت‌استرپ^۱ را برای بررسی تغییرپذیری اندازه کارایی برای هر نمونه‌گیری، طراحی نمود. در این روش از تکنیک بوت‌استرپ برای نشان دادن رتبه‌بندی و حساسیت مقادیر کارایی حاصل از روش تحلیل پوششی داده‌ها، نسبت به تغییرات ترکیب نمونه استفاده می‌شود. بوت‌استرپ یک تکنیک بازنمونه‌گیری^۲ است که توسط افرون و تیشیرانی [۱۸] ارایه شده و برای تخمین خواص توزیع نمونه‌گیری یک تخمین‌زننده (وقتی به دست آوردن آن از طریق روش‌های دیگر مشکل باشد) به کار برده می‌شود. در ساده‌ترین شکل بوت‌استرپ انتخاب تصادفی هزاران نمونه ساختگی^۳ با استفاده از نمونه‌گیری ساده تصادفی همراه با جایگذاری از مجموعه داده‌های نمونه مشاهده شده (نمونه اصلی) است، که با استفاده از هر یک از نمونه‌های ساختگی یک تخمین ساختگی (مقدار کارایی ساختگی) را به دست می‌آوریم [۱۷ و ۱۸]. این هزاران تخمین ساختگی، یک توزیع تجربی را برای تخمین‌زننده تشکیل می‌دهند و از آن به‌عنوان تخمینی از توزیع نمونه‌گیری جامعه اصلی استفاده می‌کنند. با توجه به آنچه که ذکر شد، در ادامه ابتدا به توضیح روند کلی تولید داده و بوت‌استرپ پرداخته و سپس در قسمت دوم این بخش به معرفی روش ارزیابی میزان کارایی، یعنی تحلیل پوششی داده‌ها پرداخته، و در نهایت این بخش با معرفی الگوریتم‌هایی برای تخمین روند تولید داده^۴ (DGP) در قسمت سوم، پایان خواهد یافت.

۳-۱ روند کلی تولید داده (DGP) و بوت‌استرپ

فعالیت یک واحد تولیدی که با p نهاده $(x \in R_+^p)$ ، q ستانده $(y \in R_+^q)$ تولید می‌کند را می‌توان با مجموعه تولید^۵ Ψ ، که شامل مجموعه امکان‌پذیر و فیزیکی (x, y) است، به صورت رابطه (۱) نشان داد.

$$\Psi = \{(x, y) \in R_+^{p+q} \mid x \rightarrow y\} \quad \text{نهاده } x \text{ می‌تواند ستانده } y \text{ تولید کند.} \quad (1)$$

این مجموعه را می‌توان به صورت دو بخش تشکیل دهنده آن که شامل مجموعه نهاده^۶ و مجموعه ستانده^۷ است، نشان داد؛ به طوری که مجموعه نهاده‌های تولید را برای $\forall y \in \Psi$ به صورت رابطه (۲) و همچنین مجموعه ستانده‌ها را برای $\forall x \in \Psi$ ، به وسیله رابطه (۳) نشان داد.

$$X(y) = \{x \in R_+^p \mid (x, y) \in \Psi\} \quad (2)$$

$$Y(x) = \{y \in R_+^q \mid (x, y) \in \Psi\} \quad (3)$$

رابطه بین دو مجموعه (۲) و (۳) را می‌توان به وسیله مجموعه فروض ارایه شده توسط شفارد [۱۵] که شامل فرض تحذب^۸ $X(y)$ برای تمام y ها و فرض تصرف^۹ نهاده‌ها و ستانده‌ها و ... هستند، توضیح داد. مرز کارای بیان شده

¹ Bootstrap DEA approach.

² Resampling.

³ Pseudo Sample.

⁴ Data Generation Process.

⁵ Production set.

⁶ Input set.

⁷ Output set.

⁸ Convexity.

⁹ Disposability.

بیان شده توسط فارل را می توان به صورت زیر مجموعه ای از $X(y)$ و یا $Y(x)$ ، که به ترتیب با $\partial X(y)$ و $\partial Y(x)$ نشان داده می شوند، به صورت روابط (۴) و (۵) نشان داد.

$$\partial X(y) = \{x \mid x \in X(y), \theta x \notin X(y), \forall 0 < \theta < 1\} \quad (۴)$$

$$\partial Y(x) = \{y \mid y \in Y(x), \beta y \notin Y(x), \forall \beta > 1\} \quad (۵)$$

روابطی که تابه حال بیان شدند، را می توان برای تعریف شاخص های محاسبه کارایی نهاده محور^۱ و ستانده محور^۲ محور^۳ برای بنگاه k ام (x_k, y_k) ، به صورت روابط (۶) و (۷) مورد استفاده قرار داد.

$$\theta_k = \text{Min}\{\theta \mid \theta x_k \in X(y_k)\} \quad (۶)$$

$$\beta_k = \text{Min}\{\beta \mid \beta y_k \in Y(x_k)\} \quad (۷)$$

در ادامه روش تحقیق تنها به توضیح مدل نهاده محور می پردازیم^۳. در صورتی که $\theta_k = 1$ باشد، بنگاه (x_k, y_k) به صورت کارآمد فعالیت می نماید؛ اما در صورتی که مقدار کارایی کوچک تر از یک باشد ($\theta_k < 1$)، واحد این تولیدی (x_k, y_k) به صورت ناکارآمد فعالیت نموده و می تواند مقدار y_k را با مقادیر کمتری نهاده، تولید نماید. برای تعاریف بعدی بسیار مفید خواهد بود که سطح فعالیت کارای نهاده محور بنگاه k ام را در سطح تولید y_k ، به صورت رابطه زیر تعریف نماییم:

$$X^\circ(x_k \mid y_k) = \theta x_k \quad (۸)$$

باید به این نکته توجه نمود که $X^\circ(x_k \mid y_k)$ محل نقطه تقاطع مرز کارا $\partial X(y)$ و شعاع θx_k بوده و برای محاسبه کارایی نهاده محور، نسبت فاصله شعاعی بنگاه (x_k, y_k) را به نقطه معادل آن بنگاه روی مرز کارایی $\partial X(y)$ $(X^\circ(x_k \mid y_k))$ اندازه گیری می کنیم. با توجه به اینکه مجموعه تولید (Ψ) و در نتیجه مجموعه نهاده ها $(X(y))$ و مرز کارایی تولید $(\partial X(y))$ جامعه ناشناخته هستند؛ بنابراین مقدار کارایی (θ_k) واحد تولید k ام (x_k, y_k) نیز ناشناخته خواهد بود. فرض کنید که با روش تولید داده (DGP) ، که با ρ نشان داده می شود، بتوانیم یک مجموعه نمونه تصادفی $\chi = \{(x_i, y_i) \mid i = 1, 2, \dots, n\}$ ساخته و با استفاده از آن به وسیله روش M ، برآوردی از مجموعه های مورد نظر، به شکل $\hat{\Psi}$ ، $\hat{X}(y)$ و $\partial \hat{X}(y)$ ، به دست آوریم؛ بنابراین ما می توانیم که کارایی واحد تولیدی (x_k, y_k) را به وسیله رابطه زیر تخمین بزنیم:

$$\hat{\theta}_k = \text{Min}\{\theta \mid \theta x_k \in \partial \hat{X}(y_k)\} \quad (۹)$$

باید توجه کرد که خواص نمونه ای $\hat{\Psi}$ ^۴، $\hat{X}(y)$ ، $\partial \hat{X}(y)$ و به تبع آن $\hat{\theta}_k$ تماماً بستگی به روش تولید داده ρ که روشی ناشناخته است، دارند. علاوه بر این؛ حتی اگر ρ مشخص بود، به دست آوردن آن ها به روش M بسیار مشکل می باشد؛ به خصوص هنگامی که M روشی ناپارامتریک است. در شرایطی مانند شرایط ما که مشخص کردن جزئی و تحلیلی خواص نمونه ای تخمین زن ها بسیار مشکل و یا غیرممکن است، روش بوت استرپ ممکن است که مناسب ترین و کاربردی ترین روش تخمین و تحلیل آن باشد. فرض کنید که ما با روش تولید داده

^۱ Input Oriented.

^۲ Output Oriented.

^۳ به سادگی می توان مطالبی را که برای مدل نهاده محور توضیح داده شده است، برای مدل ستانده محور نیز باز نویسی نمود.

^۴ Sampling Properties.

ρ آشنایی داشته و می‌توانیم به‌وسیله نمونه اصلی (χ) ، به یک تخمین قابل قبولی از آن مانند $\hat{\rho}$ ، دست یافته و از $\hat{\rho}$ برای تولید مجموعه داده $\chi^* = \{(x_i^*, y_i^*) | i = 1, 2, \dots, n\}$ استفاده کنیم. از این نمونه ساختگی، با روش M ، می‌توان مجموعه‌های $\hat{\Psi}^*$ ، $\hat{X}^*(y)$ ، $\partial \hat{X}^*(y_k)$ متناظر با نمونه ساختگی χ^* را تعریف کرده و در نتیجه مقدار کارایی $\hat{\theta}_k^*$ بنگاه تحت بررسی k ام (x_k, y_k) را به‌صورت رابطه (۱۰) به دست آورد [۱۹ و ۲۰].

$$\hat{\theta}_k^* = \text{Min} \{ \theta | \theta x_k \in \partial \hat{X}^*(y_k) \} \quad (10)$$

باید توجه داشت که تنها در شرایطی می‌توان از نمونه اصلی χ ، توزیع نمونه‌ای تخمین‌زن‌های $\hat{\Psi}^*$ ، $\hat{X}^*(y)$ و $\partial \hat{X}^*(y_k)$ را به‌طور کامل شناخت، که $\hat{\rho}$ شناخته‌شده باشد؛ در این صورت نیز ممکن است که محاسبه تحلیلی آن‌ها مشکل باشد؛ اما با روش مونت کارلو به راحتی می‌توان، تقریبی از توزیع‌های نمونه‌ای به دست آورد. به‌وسیله $\hat{\rho}$ ، تعداد B نمونه ساختگی $(b = 1, 2, \dots, B)$ ، χ_b^* را تولید کرده و سپس به‌وسیله روش M ، برای هر کدام از نمونه‌های ساختگی، تخمین‌های ساختگی $\hat{\Psi}_b^*$ ، $\hat{X}_b^*(y)$ ، $\partial \hat{X}_b^*(y)$ ، $(b = 1, 2, \dots, B)$ را مشخص کرده و در نهایت برای هر واحد تحت بررسی (x_k, y_k) مقادیر کارایی $\{\hat{\theta}_{k,b}^*\}_{b=1}^B$ را محاسبه می‌نماییم. تابع چگالی تجربی $\{\hat{\theta}_{k,b}^*\}_{b=1}^B$ ، تقریب مونت کارلو از توزیع $\hat{\theta}_k^*$ ، به شرط $\hat{\rho}$ است. روش بوت‌استرپ بر این ایده استوار است که در صورتی که $\hat{\rho}$ تقریب قابل قبولی از ρ باشد، توزیع شناخته‌شده بوت‌استرپ خواهد توانست که توزیع نمونه‌ای تخمین‌زن‌هایی از Ψ ، $X(y)$ ، $\partial X(y)$ ، که ناشناخته و مورد علاقه ما هستند را شبیه‌سازی کند. پس برای اندازه‌گیری مقدار کارایی θ_k بنگاه (x_k, y_k) باید رابطه (۱۱) برقرار باشد.

$$(\hat{\theta}_k^* - \hat{\theta}_k) | \hat{\rho} \sim (\hat{\theta}_k - \theta_k) | \rho \quad (11)$$

در این رابطه θ_k ، $\hat{\theta}_k$ و $\hat{\theta}_k^*$ به‌وسیله روابط به ترتیب (۶)، (۹) و (۱۰) تعریف می‌شوند. برای توضیح بیش‌تر در مورد رابطه فوق باید گفت که در صورتی رابطه فوق معتبر و صحیح خواهد بود که $\hat{\rho}$ تخمینی سازگار از ρ باشد. با توجه به رابطه (۱۱) می‌توانیم میزان تورش $\hat{\theta}_k$ را از تخمین‌زن اصلی جامعه θ_k را به‌صورت رابطه (۱۲)، به دست آوریم:

$$\text{bias}_{\rho,k} = E(\hat{\theta}_k) - \theta_k \quad (12)$$

معادل رابطه فوق در فضای بوت‌استرپ را می‌توان به‌صورت رابطه زیر بیان نمود:

$$\text{bias}_{\hat{\rho},k} = E(\hat{\theta}_k^*) - \hat{\theta}_k \quad (13)$$

مقدار مورد انتظار برای $\hat{\theta}_k^*$ را می‌توان با تقریب مونت کارلو آن به شکل رابطه زیر جایگزین کرد:

$$E(\hat{\theta}_k^*) = \frac{1}{B} \sum_{b=1}^B \hat{\theta}_{k,b}^* = \bar{\theta}_k^* \quad (14)$$

بنابراین

^۱ در این قسمت مانند مرحله چهارم الگوریتم SW عمل می‌شود. همانطور که در قسمت سوم روش تحقیق خواهید دید، در مرحله چهارم الگوریتم LT ما باید مقدار کارایی ساختگی

بنگاه ساختگی (x_k^*, y_k^*) را برای برای مشخص کردن خواص توزیع $\{\hat{\theta}_{k,b}^*\}_{b=1}^B$ مورد استفاده قرار دهیم.

$$\hat{bias}_k = \frac{1}{B} \sum_{b=1}^B \hat{\theta}_{k,b}^* - \hat{\theta}_k = \bar{\theta}_k^* - \hat{\theta}_k \quad (15)$$

با توجه به رابطه (۱۱) می توان برآوردی از (۱۲) را با استفاده از (۱۵) به دست آورد؛ بنابراین

$$\hat{bias}_{\rho,k} \approx \frac{1}{B} \sum_{b=1}^B \hat{\theta}_{k,b}^* - \hat{\theta}_k = \bar{\theta}_k^* - \hat{\theta}_k \quad (16)$$

با اصلاح مقدار تورش از برآوردگر اصلی $(\hat{\theta}_k)$ ، برآوردگر تورش-اصلاح شده به صورت رابطه (۱۷) حاصل می شود.

$$\tilde{\theta}_k = \hat{\theta}_k - \hat{bias}_{\rho,k} \approx \hat{\theta}_k - \left(\frac{1}{B} \sum_{b=1}^B \hat{\theta}_{k,b}^* - \hat{\theta}_k \right) \quad (17)$$

$$\tilde{\theta}_k \approx \hat{\theta}_k - \frac{1}{B} \sum_{b=1}^B \hat{\theta}_{k,b}^* \quad (17)$$

$$\tilde{\theta}_k \approx \hat{\theta}_k - \bar{\theta}_k^* \quad (17)$$

عبارت برآوردگر تورش-اصلاح شده را از این جهت به کار می گیریم که مقدار تورش به دست آمده برای این برآوردگر مقدار دقیق آن نیست؛ بلکه یک مقدار تقریبی از آن است. از این رو مقدار تورش برآوردگر از بین نمی رود و تنها اصلاح می شود. انحراف معیار برآوردگر $\hat{\theta}_k^*$ نیز به صورت رابطه (۱۸) نشان داده می شود [۱۹].

$$se_{\rho,k} = \left\{ \frac{1}{B-1} \sum_{b=1}^B (\hat{\theta}_{k,b}^* - \bar{\theta}_k^*)^2 \right\}^{\frac{1}{2}} \quad (18)$$

در پایان این قسمت، پس از اصلاح تورش، فاصله اطمینان θ_k را تعیین کرده و تابع توزیع تجربی $\hat{\theta}_{k,b}^*$ ، $(b=1,2,\dots,B)$ ، ارایه خواهد شد. ما به یک تابع توزیع چگالی تجربی اصلاح شده به مرکزیت تخمین زن تورش اصلاح شده $(\tilde{\theta}_k)$ از θ_k نیاز داریم؛ بنابراین تابع چگالی تجربی $\hat{\theta}_{k,b}^*$ باید به اندازه $\hat{bias}_k$ به سمت چپ منتقل می کنیم. همان طور که شکل زیر نشان می دهد، دلیل آن این است که در صورتی که به اندازه $\hat{bias}_k$ به سمت چپ منتقل شود، تابع چگالی تجربی به مرکزیت $\hat{\theta}_k$ خواهد بود تا $\tilde{\theta}_k$.

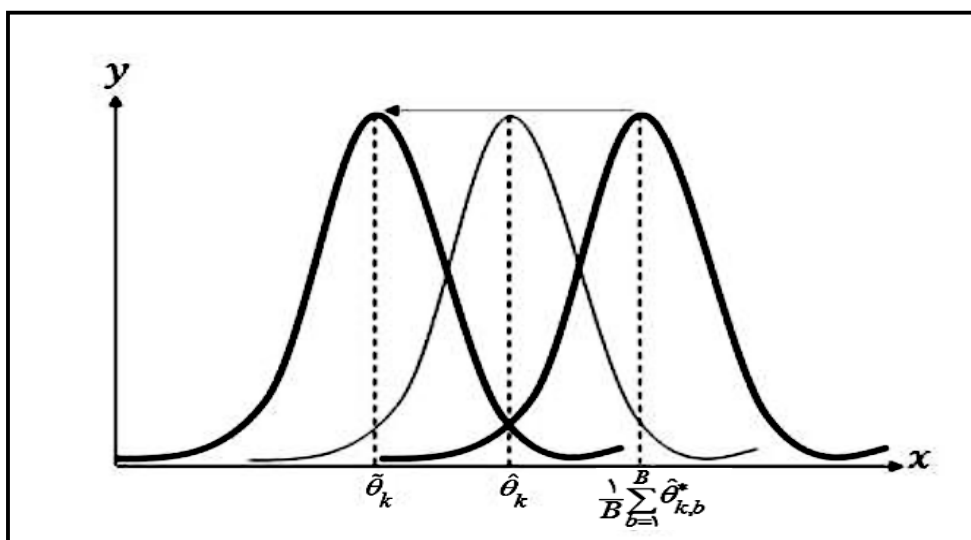
بنابراین می توان تابع چگالی تجربی $\hat{\theta}_{k,b}^*$ و سپس فاصله اطمینان θ_k را با سطح پوشش $(1-2\alpha)$ را به صورت رابطه به ترتیب (۱۹) و (۲۰) نشان داد.

$$\tilde{\theta}_{k,b}^* = \hat{\theta}_{k,b}^* - \hat{bias}_k \quad (19)$$

$$(\hat{\theta}_{k,low}, \hat{\theta}_{k,up}) = (\tilde{\theta}_k^{*(\alpha)}, \tilde{\theta}_k^{*(1-\alpha)}) \quad (20)$$

در رابطه فوق و عبارت $\tilde{\theta}_k^{*(\alpha)}$ ، مقدار بحرانی α برای تعیین فاصله اطمینان و در نتیجه تابع چگالی تجربی $\tilde{\theta}_{k,b}^*$ $(b=1,2,\dots,B)$ به کار گرفته شده است. در صورتی که تابع چگالی تجربی $\tilde{\theta}_k^{*(\alpha)}$ نامتوازن باشد، ممکن است ترجیح داده شود که میانه به عنوان مرکزیت توزیع $\tilde{\theta}_k$ انتخاب شود. سؤالی که در این قسمت پاسخ داده نشد آن است که چگونه باید $\hat{\rho}$ انتخاب شود. از آنجایی که پاسخ به این سؤال به روش تخمین M بستگی دارد، در

قسمت دوم توضیحی مختصر راجع به روش تحلیل پوششی داده‌ها (روش M) ارائه شده و پس از آن در قسمت سوم به انواع الگوریتم‌های انتخاب $\hat{\rho}$ خواهیم پرداخت [۱۹ و ۲۰].



شکل ۴. اصلاح نورش تابع توزیع [۱۶].

۳-۲ تحلیل پوششی داده‌ها

رهیافت تحلیل پوششی داده‌های CCR، کارایی فنی واحد تصمیم ساز (x_i, y_i) ، را بر اساس تعیین مجموعه تولید حاصل از نمونه $\chi = ((x_i, y_i), i = 1, 2, \dots, n)$ ، اندازه‌گیری کرده و شکل کلی آن به صورت رابطه (۲۱) نشان داده می‌شود.

$$\hat{\Psi}_{DEA} = \{(x_i, y_i) \in R^{p+q} \mid y_i \leq Yz, x_i \geq Xz, z \geq 0, i = 1, 2, \dots, n\} \quad (21)$$

برای تفکیک کارایی فنی به کارایی مقیاس و مدیریتی با اضافه کردن قید $\sum_{i=1}^n z_i = 1$ به مدل فوق، مدل BBC حاصل می‌شود. با توجه به رابطه فوق مجموعه نهاده برای سطح تولید y به صورت رابطه زیر تخمین زده می‌شود:

$$\hat{X}(y) = \{x \in R_+^p \mid (x, y) \in \hat{\Psi}_{DEA}\} \quad (22)$$

مرز کارای نهاده محور تخمین زده مربوط به مجموعه نهاده فوق، در سطح تولید ستانده y ، با $\partial \hat{X}(y)$ نشان داده می‌شود؛ این مرز زیرمجموعه‌ای از $\hat{X}(y)$ و از تعریف $\hat{\Psi}_{DEA}$ به دست می‌آید. با توجه به تعاریف فوق، برای هر واحد تصمیم گیرنده (x_i, y_i) ، مقدار کارایی تخمین زده شده $(\hat{\theta}_i)$ واحد تصمیم ساز i ام را با فرض بازدهی متغیر نسبت به مقیاس، می‌توان با حل رابطه (۲۳) به وسیله برنامه‌ریزی خطی محاسبه نمود.

$$\hat{\theta}_i = \text{Min} \left\{ \theta \mid y_i \leq Yz, \theta x_i \geq Xz; \sum_{i=1}^n z_i = 1, z > 0, 0 \leq \theta_i \leq 1, i = 1, 2, \dots, n \right\} \quad (23)$$

در حقیقت، برای تخمین میزان کارایی بنگاه i ام $(\hat{\theta}_i)$ ، نسبت فاصله شعاعی بین نقاط محل فعالیت واحد تصمیم گیرنده (x_i, y_i) و نقطه معادل آن روی مرز کارای تولید نهاده محور $(\hat{x}^\theta(x_i \mid y_i), y_i)$ ، به صورت

رابطه فوق، محاسبه می‌شود. در این رابطه $\hat{x}^\theta(x_i | y_i)$ سطحی از نهاده مصرفی است که واحد تصمیم‌ساز با سطح تولید مشخص y_i ، برای تولید کارآمد باید به آن دست یابد (با حرکت از x_i به $\hat{x}^\theta(x_i | y_i)$ در طول شعاع (θx_i) [۲۱-۲۳].

$$\hat{x}^\theta(x_i | y_i) = \hat{\theta}_i x_i \quad (24)$$

۳-۳ تخمین‌زن تحلیل پوششی داده‌های بوت‌استرپ

روش بوت‌استرپ یک روش بازنمونه‌گیری آماری است، که کارآمدی آن برای انجام استنباط آماری در مسایل پیچیده ثابت شده است. مهم‌ترین مرحله در کاربرد روش بوت‌استرپ، مشخص کردن صحیح روش ρ و یا روند تولید داده (DGP)، از داده‌های نمونه‌گیری شده از جامعه است. ایده اصلی روش بوت‌استرپ، برآورد توزیع نمونه‌ای تخمین‌زن به وسیله استفاده از توزیع تجربی تخمین‌هایی است که از بازنمونه‌گیری به دست آمده‌اند. الگوریتم‌های تحلیل پوششی داده‌های بوت‌استرپ به وسیله سیمار [۱۷]، سیمار و ویلسون (SW) [۱۹ و ۲۰] و همچنین لوسگرن و تامبور (LT) [۲۱-۲۳] بر پایه مدلی یکسان از روند تولید داده‌ها (DGP) معرفی شدند و در این مدل‌ها فرض شده است که برای تولید داده، در سطح مشخصی از ستانده‌ها، میزان نهاده برای ساخت نمونه ساختگی، از انحرافات شعاعی تصادفی از منحنی همسان مجموعه نهاده به دست می‌آید. هر نهاده در نمونه مشاهده‌شده نهاده - ستانده $\chi = ((x_i, y_i), i = 1, 2, \dots, n)$ به صورت رابطه (۲۵) نشان داده می‌شود.

$$(x_i, y_i) = \left(\frac{x^\theta(x_i | y_i)}{\theta_i}, y_i \right) \quad (25)$$

در رابطه فوق $x^\theta(x_i | y_i) \in \partial X(y_i)$ یک نقطه مشاهده نشده (روی مرز کارایی ساخته‌شده) معادل شعاعی بنگاه (x_i, y_i) ، بر روی مرز کارا $\partial X(y)$ است، که برای محاسبه کارایی مکان هندسی بنگاه را با آن نقطه مقایسه می‌کنند. مفروض است که مقادیر کارایی حقیقی از یک توزیع مشابه، مانند $\theta_i \sim F_\theta, i = 1, 2, \dots, n$. گرفته شده‌اند؛ با توجه به این نکته، می‌توان بیان کرد که مدل روند تولید داده (DGP)، ایده‌ای مشروط بر نسبت‌های ستانده‌ها و نهاده‌هاست و عناصر تصادفی روند تولید را می‌توان به طور کامل با شاخص کارایی نهاده محور تصادفی، نشان داد. ایده اصلی شبیه‌سازی بوت‌استرپ، تقلید از روند تولید داده‌ها است. الگوریتم‌های بوت‌استرپ LT و SW در هر بازنمونه به صورت زیر هستند: مشروط بر نسبت‌های نهاده‌ها و ستانده‌های مشاهده‌شده، داده‌های بازنمونه در دو گام تولید می‌شوند. در گام اول منحنی مرزی نهاده را با استفاده از نمونه مشاهده‌شده تخمین زده می‌شود؛ سپس با استفاده از مرز نهاده‌ها و مقادیر کارایی ساخته‌شده از برخی تخمین‌های توزیع F_θ ، مقادیر نهاده‌های ساختگی بوت‌استرپ به وسیله تکرار DGP معرفی شده در رابطه (۲۵)، تولید می‌شوند. این گام در الگوریتم ارایه‌شده توسط LT بر مبنای یک بازنمونه‌گیری ساده از توزیع تجربی مقادیر کارایی تخمین زده شده که در تولید کارایی‌های ساختگی مورد استفاده قرار گرفته‌اند، بنا شده است. برخلاف روش LT، الگوریتم SW در گام نخست، از فرایند بازنمونه‌گیری هموارشده استفاده کرده و این الگوریتم بر پایه استدلال سازگاری بنا شده است. در گام دوم، تخمین کارایی بوت‌استرپ به وسیله محاسبه فاصله

شعاعی مرز کارای تولیدشده به‌وسیله نمونه ساختگی، از نهاده ساختگی (در الگوریتم LT) و یا مقادیر نهاده اصلی (در الگوریتم SW) انجام می‌شود. در قسمت بعد به شرح الگوریتم‌های LT و SW و تفاوت‌های آن‌ها می‌پردازیم [۲۳].

۳-۱ الگوریتم LT

مراحل الگوریتم بوت‌استرپ ارایه‌شده توسط LT به‌صورت زیر است:

۱- با استفاده از مقادیر کارایی تخمین زده شده اصلی $\{\hat{\theta}_i, i = 1, 2, \dots, n\}$ ، بردارهای نهاده-ستانده را به‌صورت رابطه (۲۶) تبدیل می‌کنیم.

$$(\hat{x}^{\theta}(x_i | y_i), y_i) = (x_i, \hat{\theta}_i, y_i) \quad (26)$$

۲- از n کارایی فنی مجموعه تخمین زده شده $\{\hat{\theta}_{in}\}$ ، به‌صورت مستقل و همراه با جایگذاری، بازنمونه‌گیری انجام می‌دهیم. کارایی‌های به‌دست آمده از بازنمونه‌گیری با $\delta_i^*, i = 1, 2, \dots, n$ نشان داده می‌شوند.

۳- داده‌های ساختگی بوت‌استرپ به‌صورت رابطه (۲۷) تولید می‌شوند.

$$(x_i^*, y_i^*) = \left(\frac{\hat{x}^{\theta}(x_i | y_i)}{\delta_i^*}, y_i \right) \quad (27)$$

۴- تخمین زدن مقادیر کارایی بوت‌استرپ با استفاده از داده‌های ساختگی به‌وسیله حل مدل (۲۸) با برنامه‌ریزی خطی به دست می‌آید.

$$\hat{\theta}_{i,b}^{LT*} = \min_{\theta, z} \left\{ \theta | y_i \leq Yz, \theta x_i^* \geq X^* z, \sum_{i=1}^n z_i = 1, z > 0 \right\} \quad (28)$$

۵- مراحل دوم تا چهارم این الگوریتم را برای ساخت یک مجموعه B تایی کارایی بوت‌استرپ $(\hat{\theta}_{i,b}^{LT*}, b = 1, \dots, B)$ بنگاه موردنظر، B بار تکرار می‌کنیم.

در این الگوریتم مرز تخمین خورده بوت‌استرپ و مقادیر کارایی بوت‌استرپ، بر اساس بازنمونه‌گیری مقادیر کارایی فنی حاصل از توزیع تجربی مقادیر کارایی به‌دست آمده از نمونه اصلی، بازنمونه‌گیری می‌شوند. علاوه بر این، تکرارهای بوت‌استرپ کارایی بر اساس داده‌های بازنمونه‌گیری شده به همان ترتیبی که تخمین‌های اصلی بر اساس داده‌های اصلی هستند، می‌باشند [۲۲ و ۲۳].

۳-۲ الگوریتم SW

الگوریتمی که سیمار و ویلسون [۱۹ و ۲۰] (SW) معرفی کردند، در مراحل دوم و چهارم از الگوریتم لوسگرن و تامبور متفاوت است. در مرحله دوم فرایند هموارسازی، بر اساس یک هموارسازی هسته توزیع تجربی مقادیر اصلی کارایی تخمین زده شده، برای تولید یک بازنمونه هموارشده مقادیر کارایی ساختگی مورد استفاده قرار می‌گیرد. استفاده از روند هموارسازی بر اساس اصلاح بازتابی تابع چگالی هسته گوسی^۱ تخمین که توسط

¹ Gaussian.

سیلورمن [۲۴] بحث شده است، می‌باشد. این روند به صورت جزئی توسط سیمار و ویلسون بحث شده است، که در اینجا به صورت اختصار و تنها به مفاهیم اساسی آن پرداخته می‌شود. در نظر بگیرید که δ_i^* یک باز نمونه ناهموار گرفته شده به صورت مستقل با جایگذاری از توزیع تجربی تخمین‌های مقادیر اصلی کارایی فنی $\{\hat{\theta}_{in}\}$ است، روند هموارسازی در دو گام شکل می‌گیرد. اول یک اخلاص کوچک به δ_i^* اضافه شده و سپس تصحیح در توالی باز نمونه برداری اعمال می‌شود [۱۹ و ۲۴].

ابتدا یک اخلاص کوچک به مقدار $h\varepsilon_i^*$ که در آن h نشان‌دهنده پهنای باند^۱ و ε_i^* ، از یک توزیع نرمال استاندارد مستقل و یکسان (*i.i.d*) گرفته شده است، برای تولید کارایی ساختگی هموار شده $\tilde{\delta}_i^*$ ، به δ_i^* اضافه می‌شود. با توجه به این واقعیت که مقادیر کارایی در فاصله واحد محدود شده‌اند (مقادیر کارایی نهاده محور حاصل از تحلیل پوششی داده‌ها) روند بازتابی به صورت رابطه (۲۹)، در تولید $\tilde{\delta}_i^*$ استفاده شده است.

$$\tilde{\delta}_i^* = \begin{cases} \delta_i^* + h\varepsilon_i^*, & \text{if } (\delta_i^* + h\varepsilon_i^*) \leq 1 \\ 2 - (\delta_i^* + h\varepsilon_i^*), & \text{otherwise.} \end{cases} \quad (29)$$

در صورتی که مشخص شود که $\delta_i^* + h\varepsilon_i^* > 1$ است، $\tilde{\delta}_i^*$ به تصویری متقارن از بازتاب $\delta_i^* + h\varepsilon_i^*$ در نقطه مرزی تغییر می‌کند ($\tilde{\delta}_i^* = 2 - (\delta_i^* + h\varepsilon_i^*)$). یکی از مهم‌ترین مسایل در استفاده از روند هموارسازی، انتخاب پارامتر پهنای باند است (h). همان‌طور که سیلورمن [۲۴]، در مطالعه خود نشان داده است، چندین رهیافت برای انتخاب پهنای باند (h) وجود دارد. لوسگرن [۲۲] در مطالعه خود به وسیله شبیه‌سازی مونت کارلو نشان می‌دهد، که می‌توان با استفاده از یک قانون انتخاب پهنای باند خود کار و قوی برای یک متغیر که به وسیله سیلورمن ارایه شده است، مقدار h را به صورت رابطه (۳۰) محاسبه نمود [۱۹].

$$h = . / 90 \cdot n^{-\frac{1}{5}} \text{Min} \left\{ \hat{\sigma}_{\hat{\theta}}, \frac{R_{13}}{1/34} \right\} \quad (30)$$

در رابطه فوق $\hat{\sigma}_{\hat{\theta}}$ نشان‌دهنده تخمین انحراف معیار داخلی مقادیر کارایی تخمین زده شده $\{\hat{\theta}_{in}\}$ و R_{13} دامنه میان چارکی توزیع تجربی $\{\hat{\theta}_{in}\}$ است. پس از طی کردن مراحل فوق، مقادیر کارایی باز نمونه گیری هموار شده (γ_i^*) ، به وسیله تصحیح $\tilde{\delta}_i^*$ ، به صورت رابطه (۳۱) به دست می‌آیند.

$$\gamma_i^* = \frac{\bar{\delta}_i^* + (\tilde{\delta}_i^* - \bar{\delta}_i^*)}{\sqrt{1+h^2} \cdot \hat{\sigma}_{\hat{\theta}}} \quad (31)$$

در رابطه (۳۱) $\bar{\delta}_i^* = \sum_{i=1}^n \frac{\delta_i^*}{n}$ میانگین باز نمونه گیری مقادیر کارایی‌های اصلی است. دومین تفاوت بنیادین بین دو الگوریتم ارایه شده توسط LW و SW در مرحله چهارم است. در الگوریتم SW مقادیر تخمین زده شده کارایی نهاده محور بوت استرپ $\hat{\lambda}$ امین بنگاه تولیدی، با توجه به نسبت فاصله شعاعی (در سطح تولید ثابت) مقدار نهاده مصرفی $\hat{\lambda}$ امین بنگاه، که از داده‌های اصلی حاصل شده است، بر نقطه مرتبط آن، که بر روی منحنی همسان تولید

¹ Bandwidth.

ساختگی بوت‌استرپ حاصل می‌شود، به‌دست می‌آید؛ و این برخلاف الگوریتم LT که در آن از نسبت داده‌های ساختگی به مرز ساختگی بوت‌استرپ، مقادیر کارایی ساختگی حاصل می‌شوند، است [۲۵، ۲۶].

به‌طور خلاصه الگوریتم SW را می‌توان به‌صورت مراحل زیر شرح داد:

۱- بردارهای نهاده-ستانده که در محاسبه مقادیر کارایی اصلی (محاسبه شده از نمونه اصلی) $\{\hat{\theta}_i, i = 1, 2, \dots, n\}$ به‌صورت رابطه (۳۲) تبدیل می‌کنیم.

$$(\hat{x}^{\circ}(x_i, y_i), y_i) = (x_i, \hat{\theta}_i, y_i) \quad (32)$$

۲- در این مرحله به‌صورت زیر، مقادیر بازنمونه‌گیری هموار شده کارایی (γ_i^*) ساخته می‌شوند.

۱-۲- با استفاده از رابطه (۳۰) از مجموعه مقادیر کارایی‌های تخمین زده شده، برای مشخص کردن مقدار پهنای باند (h)، استفاده می‌شود.

۲-۲- مقادیر $\{\delta_i^*\}$ به‌وسیله بازنمونه‌گیری همراه با جایگذاری از توزیع تجربی تخمین زده شده مقادیر کارایی $\{\hat{\theta}_{in}\}$ تولید می‌شوند.

۲-۳- به‌وسیله رابطه (۲۹) رشته‌ای از $\{\tilde{\delta}_i^*\}$ تولید می‌شوند.

۲-۴- به‌وسیله رابطه (۳۱) مقادیر ساختگی کارایی هموار شده تولید می‌شوند.

۳- داده‌های ساختگی بوت‌استرپ به‌صورت رابطه (۳۳) تولید می‌شوند.

$$(x_i^*, y_i^*) = \left(\frac{\hat{x}^{\circ}(x_i, y_i)}{\gamma_i^*}, y_i \right) \quad (33)$$

۴- تخمین زدن مقادیر کارایی بوت‌استرپ با استفاده از داده‌های ساختگی به‌وسیله حل مدل (۳۴) با برنامه‌ریزی خطی.

$$\hat{\theta}_{i,b}^{SW*} = \min_{\theta, z} \left\{ \theta \mid y_i \leq Yz, \theta x_i \geq X^*z, \sum_{i=1}^n z_i = 1, z_i \geq 0 \right\} \quad (34)$$

۵- مراحل دوم تا چهارم این الگوریتم را برای ساخت یک مجموعه B تایی کارایی بوت‌استرپ $(\hat{\theta}_{i,b}^{SW*}, b = 1, \dots, B)$ بنگاه موردنظر، B بار تکرار می‌کنیم [۲۷].

۳-۳-۳ الگوریتم LSW

الگوریتم سوم نیز با نام LSW، بر اساس ترکیب الگوریتم LT و SW تعریف می‌شود. همان‌طور که ذکر شد به‌جز مراحل دوم و چهارم، سایر مراحل الگوریتم‌های SW و LT، مشابه هستند. بر این اساس برای آرایه الگوریتم LSW، برای مرحله دوم روند هموارسازی الگوریتم SW و برای مرحله چهارم، استفاده از داده‌های نهاده ساختگی که در تخمین‌زن کارایی بوت‌استرپ الگوریتم LT تعریف شده است، مورد استفاده قرار می‌گیرد.

۴ داده‌های تحقیق

در این مطالعه از داده‌های بخش صنعت کشور که در پایگاه اطلاعاتی مرکز آمار ایران^۱ موجود است، استفاده شده است. این اطلاعات حاصل نتایج نمونه‌گیری از کارگاه‌های صنعتی ۱۰ نفر کارکن و بیش‌تر که با مراجعه به ۱۴۶۶۴ کارگاه صنعتی در سال ۱۳۹۰ به‌دست آمده است، می‌باشد. در این پایگاه داده، داده‌ها در قالب کدهای دو، سه و چهار رقمی سومین ویرایش طبقه‌بندی استاندارد بین‌المللی فعالیت‌های اقتصادی^۲ طبقه‌بندی شده‌اند. توزیع کارگاه‌های صنعتی در سال ۱۳۹۰ نشان می‌دهد که ۲۱/۵۹ درصد در صنایع تولید سایر محصولات کانی غیرفلزی (کد ۲۶)، ۱۸/۳۳ درصد در صنایع مواد غذایی و آشامیدنی (کد ۱۵) و ۷/۸۵ درصد در صنایع تولید محصولات فلزی فابریکی به‌جز ماشین‌آلات و تجهیزات اشتغال داشته‌اند. داده‌ها مورد استفاده در این مطالعه به‌قرار زیر هستند:

- (۱) ارزش تولید: شامل مجموع ارزش کالاهای تولیدشده، ارزش ضایعات قابل فروش و تغییرات ارزش کالاهای در جریان ساخت.
 - (۲) ارزش مواد خام و اولیه: شامل ارزش موادی که به‌منظور تغییر شکل فیزیکی یا شیمیایی به کارگاه وارد و به مصرف می‌رسد است. این مواد ممکن است خام یا نیمه ساخته باشد که برای مراحل بعدی عملیات تولید کالا در کارگاه به کار گرفته می‌شوند.
 - (۳) ارزش انرژی: شامل ارزش سوخت مصرف‌شده و برق خریداری‌شده.
 - (۴) جبران خدمات مزد و حقوق‌بگیران: شامل مزد و حقوق و سایر پرداختی‌ها (پول و کالا و ...) به مزد و حقوق‌بگیران.
- در جدول (۱) خلاصه وضعیت نهاده‌های یادشده درج شده است [۱].

جدول ۱. ارزش تولید و نهاده‌های بخش صنعت کشور به ترتیب کدهای ۴ رقمی ISIC در سال ۱۳۹۰ (ریال)

تولید	موجودی سرمایه	مواد اولیه	دستمزد نیروی کار	انرژی
۵۹۵۴۵۷۱۹۸۶۸۶۳۲۰	۶۸۹۸۰۲۲۲۸۳۲۴۳	۵۵۳۰۹۶۲۳۳۲۸۹۳۲۹	۶۷۶۸۳۳۴۶۱۹۶۷۵	۵۳۳۲۹۷۹۳۳۸۱۱۴
(کد ۲۳۲۰)	(کد ۲۷۱۰)	(کد ۲۳۲۰)	(کد ۳۴۱۰)	(کد ۲۳۲۰)
۱۵۰۰۷۷۵۳۲۷۵۰۸۴	۴۹۹۵۹۲۳۶۳۴۳۱۲	۱۱۲۱۲۶۷۰۸۶۷۳۰۸	۵۶۰۱۵۱۱۴۲۷۳۸	۲۴۱۲۲۱۹۹۶۴۶۶
۱۸۹۵۰۰۰۰۰۰	۳۷۵۸۰۷۷۴۳۱۶	۱۴۴۰۵۱۰۰۰۰	۴۵۲۰۰۰۰۰۰	۸۷۷۰۰۰۰
(کد ۱۷۲۵)	(کد ۲۳۱۰)	(کد ۱۷۲۵)	(کد ۱۷۲۵)	(کد ۱۷۲۵)
۵۸۸۰۴۰۷۶۵۷۸۰۲۲	۱۰۴۴۱۷۵۵۲۴۷۱۰۰	۵۲۰۹۴۱۰۳۵۲۹۶۸۶	۹۴۸۳۱۳۴۸۳۸۶۲	۷۸۹۸۶۴۵۹۰۵۸۵

منبع: یافته‌های تحقیق

^۱ www.amar.org.ir.

^۲ IS.I.CREV.3

۵ یافته‌های تحقیق

در این مطالعه به وسیله الگوریتم LSW که بر اساس ترکیب الگوریتم‌های LT و SW تعریف می‌شود، مقادیر کارایی تورش اصلاح شده و فواصل اطمینان مرتبط با آن‌ها با استفاده از رویکرد تحلیل پوششی داده‌های نهاده محور، با ۱۰۰۰ تکرار (بازنمونه‌گیری $(B = 1000)$)، در فاصله اطمینان ۹۰ درصد و با فرض بازده ثابت (CRS^1) و متغیر (VRS^2) نسبت به مقیاس تخمین زده شد و نتایج آن در جداول ۲ و ۳ درج گردید.

جدول ۲. مقادیر کارایی فنی تورش اصلاح شده و فواصل اطمینان تخمینی بخش صنعت کشور

کد ISIC	میانگین تخمین زده شده کارایی تورش اصلاح شده انحراف معیار تخمین زده شده				
	فاصله اطمینان				
	با فرض بازدهی متغیر نسبت به مقیاس (CRS)				
	پایین			بالا	
۱۵	۰/۷۵۵۶	۰/۷۲۰۵	۰/۱۲۵۸۵۹	۰/۵۱۳۴۶۴	۰/۹۰۳۳۹۵
۱۶	۰/۷۵۲	۱	۰/۱۲۹۱	۰/۷۸۷۷	۱
۱۷	۰/۷۵۴۸۶	۰/۶۷۵۲۵	۰/۱۲۶۰۳	۰/۴۷۲۳۷	۰/۸۶۲۲۷
۱۸	۰/۷۶۰۴	۰/۷۹۱۳	۰/۱۲۷۷	۰/۵۸۱۲	۱
۱۹	۰/۷۵۹۴۳۳	۰/۶۳۸۴	۰/۱۲۷۲۳۳	۰/۴۲۹۱۶۷	۰/۸۴۷۵۶۷
۲۰	۰/۷۵۶۳۸	۰/۷۸۹۵۲	۰/۱۲۷۱	۰/۵۸۰۴۸	۰/۹۶۲۰۲
۲۱	۰/۷۵۲۴۶۷	۰/۶۱۳۵۳۳	۰/۱۲۵۸۶۷	۰/۴۰۶۵۳۳	۰/۸۲۰۵۳۳
۲۲	۰/۷۵۸۹۴	۰/۷۶۴۰۸	۰/۱۲۷۵۴	۰/۵۵۴۳	۰/۹۳۱۶۴
۲۳	۰/۷۵۶۰۵	۰/۸۳۹۰۵	۰/۱۲۳۳۵	۰/۶۳۶۱۵	۰/۹۴۰۷۵
۲۴	۰/۷۵۷۳۱۱	۰/۸۴۲۳	۰/۱۲۵۷۵۶	۰/۶۴۰۵۸۹	۰/۹۴۹۱۸۹
۲۵	۰/۷۵۲	۰/۶۱۸۳	۰/۱۲۵۸	۰/۴۱۱۴	۰/۸۲۵۲۶۷
۲۶	۰/۷۵۸۷	۰/۶۹۸۴۵	۰/۱۲۶۵۷	۰/۴۹۰۲۸	۰/۸۸۵۹۵
۲۷	۰/۷۵۷۹	۰/۷۴۳۸۶۷	۰/۱۲۸۳۸۳	۰/۵۳۲۶۵	۰/۹۳۶۴۸۳
۲۸	۰/۷۵۶۵۳۳	۰/۷۳۲۳۳۳	۰/۱۲۷۷۸۳	۰/۵۲۲۱۶۷	۰/۹۱۴۹۸۳
۲۹	۰/۷۵۴۶۴۳	۰/۷۲۴۱۹۳	۰/۱۲۵۴۵	۰/۵۱۷۸۴۳	۰/۹۰۸۹۲۹
۳۰	۰/۷۶۳۶	۱	۰/۱۲۸۱	۰/۷۸۹۲	۱
۳۱	۰/۷۵۵۵۵	۰/۸۰۸۳۱۷	۰/۱۲۵۷۳۳	۰/۶۰۱۴۸۳	۰/۹۷۳۶۵
۳۲	۰/۷۵۰۶۶۷	۰/۷۶۰۱۶۷	۰/۱۲۵۰۳۳	۰/۵۵۴۵۳۳	۰/۹۵۸۶
۳۳	۰/۷۵۷۷۲	۰/۷۸۹۳۴	۰/۱۲۶۱۶	۰/۵۸۱۸۲	۰/۹۲۴۵۸
۳۴	۰/۷۵۸۹۶۷	۰/۸۰۰۳	۰/۱۲۵۴۳۳	۰/۵۹۴۰۳۳	۰/۹۳۸۸۳۳
۳۵	۰/۷۵۹۲۵	۰/۷۲۰۲	۰/۱۲۷۲	۰/۵۱۱	۰/۹۲۹۴۵
۳۶	۰/۷۵۸۴۶	۰/۷۴۸۴	۰/۱۲۷۷۸	۰/۵۳۸۲	۰/۹۳۷۱۸
۳۷	۰/۷۵۵	۰/۱۶۸۵	۰/۱۲۴۶	.	۰/۵۷۸۴
میانگین	۰/۷۵۶۶۲۷	۰/۷۳۸۵۳۵	۰/۱۲۶۵۰۳	۰/۵۳۲۴۵۹	۰/۹۰۹۹۸۶

منبع: یافته‌های تحقیق

¹ Constant Return to Scale.

² Variable Return to Scale.

جدول ۳. مقادیر کارایی مدیریتی و مقیاس تورش اصلاح شده و فواصل اطمینان تخمینی بخش صنعت کشور

کد ISIC	تخمین میانگین	کارایی تورش اصلاح شده	انحراف معیار تخمین زده شده	فاصله اطمینان		کارایی مقیاس
				پایین	بالا	
۱۵	۰/۸۲۸۹۲۳	۰/۷۵۲۴۵	۰/۱۲۵۲۶۴	۰/۵۴۶۴۰۵	۰/۹۲۷۲۷۳	۰/۹۵۷۴۶۸
۱۶	۰/۸۳۰۱	۱	۰/۱۲۶۷	۰/۷۹۱۶	۱	۱
۱۷	۰/۸۲۸۴۱	۰/۷۱۹۱۱	۰/۱۲۴۳۵	۰/۵۰۸۶	۰/۹۰۳۶۹	۰/۹۱۶۲۰۱
۱۸	۰/۸۲۱۶	۰/۸۰۷۵	۰/۱۲۵۲	۰/۶۰۱۶	۱	۰/۹۷۹۹۳۸
۱۹	۰/۸۳۰۱۳۳	۰/۷۰۳۲	۰/۱۲۵۸۶۷	۰/۴۹۶۱۶۷	۰/۹۱۰۱۶۷	۰/۹۱۴۳۳۹
۲۰	۰/۸۲۸	۰/۸۳۵۹	۰/۱۲۴۳۶	۰/۶۳۱۳۶	۰/۹۶۶۸۲	۰/۹۵۰۴۶۲
۲۱	۰/۸۲۶۳	۰/۵۶۱۶	۰/۱۲۴۸	۰/۳۱۷۹	۰/۸۳۵۶۳۳	۰/۹۵۳۸۱۴
۲۲	۰/۸۲۹۶۴	۰/۷۹۳۱۲	۰/۱۲۴۲۴	۰/۵۸۸۷۶	۰/۹۴۸۰۴	۰/۹۶۳۷۸
۲۳	۰/۸۲۸۴	۱	۰/۱۲۶۱۵	۰/۷۹۲۵	۱	۰/۸۳۹۰۵
۲۴	۰/۸۲۸۱۸۹	۰/۸۵۲۰۵۶	۰/۱۲۵۷۸۹	۰/۶۴۵۱۴۴	۰/۹۵۵۸۵۶	۰/۹۸۶۹۸۹
۲۵	۰/۸۳۰۲۳۳	۰/۶۸۶۵۶۷	۰/۱۲۷۱	۰/۴۷۷۵	۰/۸۹۵۶	۰/۹۰۳۴۰۲
۲۶	۰/۸۳۱	۰/۷۱۹۸۴	۰/۱۲۵۹۷	۰/۵۱۲۶۶	۰/۹۰۵۴۱	۰/۹۷۰۳۳۶
۲۷	۰/۸۳۵۸۳	۰/۷۹۱۸	۰/۱۲۵۲۸۳	۰/۵۸۵۷۱۷	۰/۹۳۶۲۶۷	۰/۹۵۰۰۴۸
۲۸	۰/۸۲۹۶۵	۰/۷۶۷۹	۰/۱۲۵۰۶۷	۰/۵۶۲۱۵	۰/۹۳۵۹۳۳	۰/۹۵۵۳۴۷
۲۹	۰/۸۲۸۶۷۱	۰/۷۱۸۴۹۳	۰/۱۲۴۱۸۶	۰/۵۰۴	۰/۹۲۰۰۸۶	۰/۹۵۸۳۴۸
۳۰	۰/۸۲۴۳	۱	۰/۱۲۷۱	۰/۷۹۱	۱	۱
۳۱	۰/۸۲۹۰۶۷	۰/۸۳۰۸۱۷	۰/۱۲۴۷۵	۰/۶۲۵۶۵	۰/۹۸۳۳۵	۰/۹۷۲۸۱۷
۳۲	۰/۸۲۹۵۶۷	۰/۷۶۶۸۶۷	۰/۱۲۳۹۶۷	۰/۵۶۲۹۳۳	۰/۹۵۹۵۶۷	۰/۹۹۱۴۴
۳۳	۰/۸۲۸۲۸	۰/۸۵۵۰۲	۰/۱۲۴۵۶	۰/۶۵۰۱۸	۰/۹۶۰۲	۰/۹۲۶۷۲۶
۳۴	۰/۸۳۲۵	۰/۸۷۲۸	۰/۱۲۴۶۳۳	۰/۶۶۷۷۳۳	۰/۹۷۲۷۳۳	۰/۹۱۹۶۳۹
۳۵	۰/۸۲۸۴۳۳	۰/۷۶۰۴۶۷	۰/۱۲۶۱۳۳	۰/۵۵۳	۰/۹۴۵۶۵	۰/۹۵۴۵۱۱
۳۶	۰/۸۳۱۱۲	۰/۸۵۹۵۸	۰/۱۲۷۰۶	۰/۶۵۰۵۶	۰/۹۷۰۲۸	۰/۸۷۹۹۶۶
۳۷	۰/۸۳۳۷	۰/۳۲۱۵	۰/۱۲۸	۰	۰/۷۴۲۶	۰/۵۲۴۱۰۶
میانگین	۰/۸۲۹۱۲۲	۰/۷۸۱۵۹۱	۰/۱۲۵۵۰۱	۰/۵۶۷۹۶۲	۰/۹۳۸۰۵	۰/۹۲۹۰۷۵

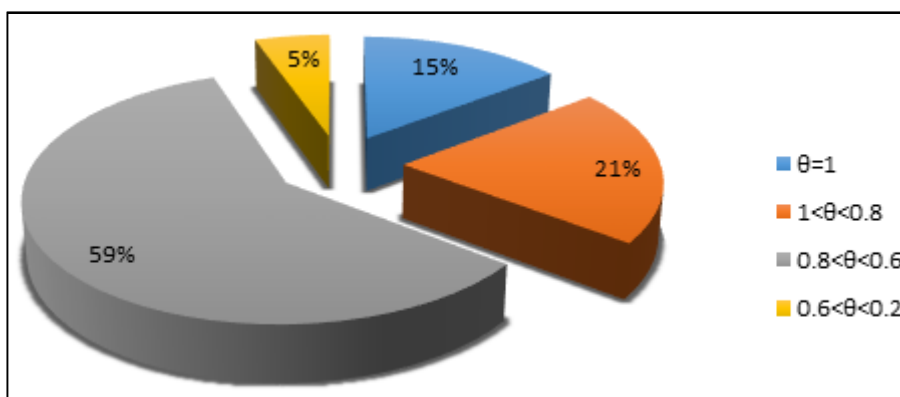
منبع: یافته‌های تحقیق

به دلیل آنکه فرض بازدهی ثابت نسبت به مقیاس تنها در مقیاس مناسب فعالیت دارای اعتبار بوده و در صورتی که فعالیت بنگاه‌ها در این مقیاس انجام نشود منجر به به دست آمدن نمرات کارایی فنی ناخالص می‌گردد، مدل دیگری نیز با فرض بازدهی متغیر نسبت به مقیاس برای تجزیه کارایی فنی به نمرات کارایی خالص (کارایی مدیریتی) و کارایی مقیاس، تخمین زده شد^۱ و در نتیجه علاوه بر به دست آوردن نمرات کارایی مدیریتی، نمرات کارایی مقیاس با استفاده از رابطه (۳۵) محاسبه گردید.

$$(۳۵) \quad \text{کارایی مقیاس} \times \text{کارایی مدیریتی} = \text{کارایی فنی}$$

^۱ مدل VRS با مقید کردن مدل CRS به شرط $\sum_{i=1}^n Z_i = 1$ به دست می‌آید.

پهنای باند مورد استفاده در این مطالعه به وسیله رابطه (۳۰) که توسط سیلورمن ارایه شده است، مورد محاسبه قرار گرفته و مقدار به دست آمده آن برای مدل‌های به ترتیب CRS و VRS برابر با ۰/۰۰۰۰۷ و ۰/۰۰۰۰۸ است. مقادیر کارایی در این جداول که به ترتیب کدهای ISIC دو رقمی گزارش شده است توسط میانگین‌گیری از نمرات کارایی محاسبه شده از داده‌های ۱۳۰ صنعت به ترتیب کدهای ISIC چهار رقمی به دست آمده است. همان‌طور که جدول شماره (۳) نشان می‌دهد کارایی مقیاس بخش صنعت کشور که از نسبت مقادیر کارایی فنی و مدیریتی حاصل می‌شود، به طور میانگین حدود ۰/۹۳ است و به جز صنایع تولید زغال کک (کد ۲۳)، بازیافت (کد ۳۷) و تولید مبلمان و مصنوعات طبقه‌بندی نشده (کد ۳۶)، مقادیر کارایی مقیاس سایر صنایع بیش از ۰/۹ می‌باشد و این نشان‌دهنده آن است که مهم‌ترین علت ناکارایی فنی بخش صنعت که ارقام آن‌ها در جدول ۲ آمده است، ناکارآمدی فنی خالص یا مدیریتی است که از این لحاظ همان‌طور که در جدول شماره (۴) نشان داده شده، صنایع تولید کاغذ و محصولات کاغذی (کد ۲۱) و صنعت بازیافت (کد ۳۷) با کارایی کمتر از ۶۰ درصد، پایین‌ترین میزان کارایی مدیریتی را به خود اختصاص داده‌اند. همچنین با در نظر گرفتن نمرات کارایی مدیریتی ۱۳۰ صنعت مورد بررسی به ترتیب کدهای چهار رقمی ISIC، همان‌طور که شکل شماره (۵) نیز گویای آن است، تقریباً ۱۵ درصد صنایع کشور به صورت کاملاً کارآمد فعالیت کرده و پنج درصد صنایع با ناکارآمدی مدیریتی بسیار پایینی فعالیت می‌کنند.



شکل ۵. بازه نمرات کارایی مدیریتی بخش صنعت کشور به ترتیب کدهای چهاررقمی

۶ نتیجه‌گیری

هدف محوری این مقاله سنجش ضریب کارایی صنایع کارخانه‌ای ایران با توجه به داده‌های بوت‌استرپ و رهیافت LSW می‌باشد. برای سنجش ضریب کارایی از اطلاعات کدهای آیسیک بخش صنعت استفاده شده است. یافته‌های تحقیق نشان می‌دهد کارایی مقیاس بخش صنعت کشور به طور میانگین ۰/۹۳ است و به جز صنایع بازیافت (کد ۳۷) و تولید مبلمان و مصنوعات طبقه‌بندی نشده (کد ۳۶)، مقادیر کارایی مقیاس سایر صنایع بیش از ۰/۹ می‌باشد و این نشان‌دهنده آن است که مهم‌ترین علت ناکارایی فنی بخش صنعت ناکارآمدی فنی خالص یا مدیریتی است. همچنین یافته‌های تحقیق مؤید آن است که تقریباً ۱۵ درصد از صنایع کشور به صورت کاملاً کارآمد فعالیت کرده و پنج درصد از صنایع با ناکارآمدی مدیریتی بسیار پایینی فعالیت می‌کنند؛ پایین‌ترین نمرات

کارایی مدیریتی را نیز صنایع به ترتیب تولید ماشین ابزارها (کد ۲۹۲۲) و تکمیل منسوجات (کد ۱۷۱۲) با نمرات کارایی مدیریتی پایین تر از ۰/۳ به خود اختصاص داده اند.

منابع

- [۱] دفتر صنعت و معدن و امور زیربنایی (۱۳۹۲). نتایج آمارگیری از کارگاه های صنعتی ۱۰ نفر کارکن و بیشتر کشور سال ۱۳۹۰. مرکز آمار ایران.
- [۲] آزادی نژاد، ع.، آماده، ح.، امامی میدی، ع.، (۱۳۹۲). بررسی عوامل مؤثر در کارایی فنی بخش صنعت استان های کشور با رویکرد تحلیل پوششی داده ها. مجله تحقیقات اقتصادی، ۴۹ (۱)، ۱۷۳-۱۸۸.
- [۱۲] زراءنژاد، م.، خداداد کاشی، ف.، یوسفی، ر.، (۱۳۹۱). ارزیابی کارایی فنی صنایع کارخانه ای ایران. فصلنامه اقتصاد مقداری، ۹ (۲)، ۳۱-۴۸.
- [۱۳] شهیکی تاش، م.، طاهرپور، ج.، شیوایی، ا.، (۱۳۹۳). ارزیابی عوامل مؤثر بر ناکارایی فنی صنایع کارخانه ای ایران (رهیافت تابع مرزی تصادفی و روش حداکثر درستمایی). فصلنامه پژوهشنامه اقتصادی، ۱۴ (۵۲)، ۲۷-۴۷.
- [۱۶] بهادری، ع.، حسینی نهاد، س.، حبیبی نیا، ق.، (۱۳۹۲). استفاده از فرایند شبیه سازی بوت استرپ برای برآورد مرز کارای ناپارامتری «بررسی مشکلات موجود در فرآیند ارایه شده در مقاله سعید عبادی». مجله تحقیق در عملیات در کاربردهای آن، ۱۰ (۲)، ۱۳۵-۱۱۳.
- [3] Halkos, G., Tzeremes, N., (2010). Performance evaluation using bootstrapping DEA techniques: Evidence from industry ratio analysis. MPRA Paper.
- [4] Lee, B. L., Worthington, A. C., (2014). Technical efficiency of mainstream airlines and low-cost carriers: New evidence using bootstrap data envelopment analysis truncated regression. Journal of Air Transport Management, 38, 15-20.
- [5] Zakaria, S., Salleh, M., (2014). A Bootstrap Data Envelopment Analysis (BDEA) approach in Islamic banking sector: A method to strengthen efficiency measurement. Industrial Engineering and Engineering Management (IEEM), 2014 IEEE International Conference on, IEEE.
- [6] De Jorge-Moreno, J., Rojas Carrasco, O., (2015). Technical efficiency and its determinants factors in Spanish textiles industry (2002-2009). Journal of Economic Studies, 42(3).
- [7] Brümmer, B., (2001). Estimating confidence intervals for technical efficiency: the case of private farms in Slovenia. European review of agricultural economics, 28(3), 285-306.
- [8] Gocht, A., Balcombe, K., (2006). Ranking efficiency units in DEA using bootstrapping an applied analysis for Slovenian farm data. Agricultural Economics, 35(2), 223-229.
- [9] Dong, F., Featherstone, A. M., (2006). Technical and scale efficiencies for chinese rural credit cooperatives: a bootstrapping approach in data envelopment analysis. Journal of Chinese Economic and Business Studies, 4(1), 57-75.
- [10] Balcombe, K., Fraser, I., (2008). An application of the DEA double bootstrap to examine sources of efficiency in Bangladesh rice farming. Applied Economics, 40(15), 1919-1925.
- [11] Odeck, J., (2009). Statistical precision of DEA and Malmquist indices: A bootstrap application to Norwegian grain producers. Omega, 37(5), 1007-1017.
- [14] Farrell, M. J., (1957). The measurement of productive efficiency. Journal of the Royal Statistical Society, Series A (General), 253-290.

- [15] Charnes, A., Cooper, W. W., Rhodes, E., (1978). Measuring the efficiency of decision making units. *European Journal of Operational Research*, 2(6), 429–444.
- [17] Simar, L., (1996). Aspects of statistical analysis in DEA-type frontier models. *Journal of productivity analysis*, 7(3), 177-185.
- [18] Efron, B., Tibshirani, R., (1993). An introduction to the bootstrap: Monographs on Statistics and Applied Probability. New York and London: Chapman and Hall/CRC.
- [19] Simar, L., Wilson, P. W., (1998). Sensitivity analysis of efficiency scores: How to bootstrap in nonparametric frontier models. *Management science*, 44(1), 49-61.
- [20] Simar, L., Wilson, P. W., (2000). A general methodology for bootstrapping in non-parametric frontier models. *Journal of applied statistics*, 27(6), 779-802.
- [21] Löthgren, M., Tambour, M., (1996). Scale Efficiency and Scale Elasticity in DEA-models: A Bootstrapping Approach. *Economic Research Inst.*
- [22] Lothgren, M., (1998). How to bootstrap DEA estimators: a Monte Carlo comparison. *WP in Economics and Finance*, 223.
- [23] Lothgren, M., Tambour, M., (1999). Testing scale efficiency in DEA models: a bootstrapping approach. *Applied Economics*, 31(10), 1231-1237.
- [24] Silverman, B. W., (1986). *Density Estimation for Statistics and Data Analysis*, Chapman and Hall Ltd, London.
- [25] Nguyen, H. O., Nguyen, H. V., Chang, Y. T., Chin, A. T. H., Tongzon, J., (2016). Measuring port efficiency using bootstrapped DEA: the case of Vietnamese ports. *Maritime Policy & Management*, 43(5), 644–659.
- [26] Stewart, C., Matousek, R., Nguyen, T. N., (2016). Efficiency in the Vietnamese banking system: A DEA double bootstrap approach. *Research in International Business and Finance*, 36, 96–111.
- [27] Wanke, P., Barros, C. P., (2016). New evidence on the determinants of efficiency at Brazilian ports: a bootstrapped DEA analysis. *International Journal of Shipping and Transport Logistics*, 8(3), 250–272.