

استفاده از فرآیند شبیه سازی بوت استرپ برای برآورد مرز تولید کارای ناپارامتری بررسی مشکلات موجود در فرآیند ارایه شده در مقاله سعید عبادی

علیرضا بهاری^{*}، سعید حسینی نهاد^۱، قاسم حبیبی نیا^۲

۱- دانشگاه آزاد اسلامی واحد قم، قم، ایران

۲- دانشگاه علم و صنعت ایران، تهران، ایران

۳- دانشگاه آزاد اسلامی واحد شهرکرد، شهرکرد، ایران

رسید مقاله: ۲۱ آبان ۱۳۹۱

پذیرش مقاله: ۲۲ فروردین ۱۳۹۲

چکیده

تحلیل پوششی داده‌ها (DEA) یک روش ناپارامتری برای محاسبه‌ی (برآورد) کارایی گروهی از واحدهاست که فعالیت یکسانی را انجام می‌دهند. مرز به دست آمده از این روش یک مرز نسبی قابل دسترس در دنیای واقعی می‌باشد. درست برخلاف روش پارامتری که در آن واحدها نسبت به مرزی سنجیده می‌شوند که عموماً در دنیای واقعی، غیر قابل دسترس است. کارایی به دست آمده از این روش یک مقدار نسبی و نه مقدار واقعی آن می‌باشد. به عبارت دیگر کارایی به دست آمده از این روش برآوردی از مقدار واقعی کارایی است. به دلیل نامشخص بودن توزیع جامعه میزان دقت کارایی برآورد شده مورد سوال قرار می‌گیرد. فرآیند شبیه‌سازی بوت استرپ که توسط افرون [۱] معرفی شده است، می‌تواند برای بالا بردن میزان دقت کارایی برآورد شده مورد استفاده قرار گیرد. در این مقاله ضمن تشریح کامل فرآیند بوت استرپ و نحوه استفاده از آن، به اشکالات موجود در مقاله‌ی سعید عبادی که بدون توجه به برخی نارسایی‌های این روش، به کارگرفته شده است؛ اشاره می‌کنیم. در این مقاله از روش بوت استرپ ساده استفاده شده که با وجود تذکراتی که افرون در استفاده از این روش و به خصوص برای مدل‌های مرزی داده بود، فریر و هیرچبرگ به گونه‌ای دیگر از این روش استفاده کردند. علاوه بر این، در مقاله سعید عبادی [۲] هیچ اشاره‌ای به پیشینه‌ی بوت استرپ نشده و خواننده گمان می‌کند که کار توسط نویسنده مقاله انجام شده است. اما این کار مربوط به مقاله فریر- هیرچبرگ [۳] می‌باشد که بعدها به دلیل برخی نارسایی‌ها این روش منسوخ شد [۴].

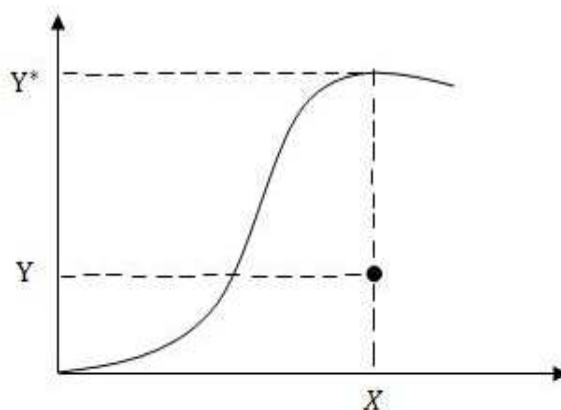
کلمات کلیدی: تحلیل پوششی داده‌ها، DEA، بوت استرپ، مرز کارای ناپارامتری.

* عهده دار مکاتبات

آدرس الکترونیک: ar_bahari@yahoo.com

۱ مقدمه

برای محاسبه‌ی کارایی واقعی نیاز داریم تا از بیشترین خروجی قابل تولید از یک ورودی اطلاع داشته باشیم. به عبارت دیگر، به یک تابع تحت عنوان تابع تولید نیاز داریم که مقدار تابع به ازای هر ورودی، نشان دهنده‌ی بیشترین میزان قابل تولید از آن ورودی باشد (شکل ۱). متأسفانه در دنیای واقعی، قادر به محاسبه بیشترین میزان تولید شده از یک ورودی مشخص نیستیم. به عبارت دیگر مرز کارایی واقعی نامشخص است.



شکل ۱. ورودی X برای تولید خروجی Y استفاده شده است و Y^* بیشترین میزان قابل تولید از ورودی X است.

برای فائق آمدن بر این مشکل و آرایه‌ی مقدار کارایی هر واحد، می‌توان تابعی را به عنوان تابع تولید مشخص کرد. به عبارت دیگر یک تخمین از مرز واقعی را به جای آن بنشانیم. این کار به دو روش پارامتری و ناپارامتری انجام می‌پذیرد.

در مدل‌های پارامتری، شکل تابعی خاصی با برخی پارامترهای مجهول برای کارایی در نظر گرفته می‌شود. در این روش هدف مشخص کردن پارامترهای مجهول به کمک داده‌های نمونه موجود می‌باشد. یکی از بزرگ‌ترین معایب روش پارامتری، تحمیل شکل تابعی خاصی به مدل است. علاوه بر آن واحدها نسبت به مرزی ارزیابی می‌شوند که معمولاً در دنیای واقعی غیر قابل دسترس می‌باشد. در روش ناپارامتری هیچ شکل تابعی خاصی به مدل تحمیل نمی‌شود و اجازه داده می‌شود تا داده‌های نمونه، خود شکل تابع تولید را مشخص کنند. در این روش به جای آنکه واحدها با یک مرز غیر قابل دسترس پارامتری ارزیابی شوند؛ با یک مرز قابل دسترس ساخته شده از داده‌ها ارزیابی می‌شوند. یکی از ضعف‌های این روش تغییر محسوس مقادیر کارایی، در صورت تغییر نمونه است که دلیل اصلی استفاده از بوت استرپ برای تخمین بهتر می‌باشد.

با توجه به آنچه گفته شد؛ در اینجا با مسأله‌ی تخمین مرز کارایی مواجه هستیم. ابتدا می‌بایست برآوردی از مرز کارایی به دست آورده؛ سپس کارایی هر واحد را نسبت به آن مرز محاسبه کنیم. واضح است که هر چه مرز برآورد شده بهتر باشد؛ مقادیر کارایی به دست آمده، به مقادیر کارایی‌های واقعی نزدیک‌تر خواهد بود. در ادامه و در بخش ۲ به بررسی برخی خصوصیات مهم مرز تحلیل پوششی داده‌ها می‌پردازیم. در بخش ۳ ساختار کلی فرآیند بوت استرپ تشریح داده شده است. بخش ۴ اختصاص به روش‌های تولید داده (DGP) دارد. در این بخش

ابتدا در مورد روش بوت استرپ ساده (فرآیند استفاده شده در مقاله سعید عبادی) صحبت می‌کنیم و سپس دلایل عدم استفاده از این روش در مدل‌های مرزی را مورد بحث و بررسی قرار می‌دهیم و در انتهای این بخش فرآیند اصلاح شده [۴] را تشریح می‌نماییم.

۲. برآورد مرز ناپارامتری

بنکر [۵] با در نظر گرفتن شکل تابعی زیر برای مرز کارایی، نشان داد که مرزی که به دست می‌آید؛ در واقع همان مرز مشخص شده به وسیله DEA است.

$$1- g(x) \text{ یکنواست. یعنی اگر } x'' \leq x' \text{ آن گاه } g(x'') \leq g(x').$$

$$2- g(x) \text{ محدب است. یعنی اگر } x' \text{ و } x'' \text{ در دامنه تابع باشند و } x = \gamma x' + (1-\gamma)x'' \text{ که}$$

$$0 < \gamma < 1 \text{ آن گاه } g(x) \geq \gamma g(x') + (1-\gamma)g(x'').$$

$$3- \text{ برای هر مشاهده } (x_j, y_j) \text{ داریم } g(x_j) \geq y_j \text{ که } j = 1, \dots, n.$$

$$4- \text{ برای هر تابع } g(x) \text{ که در شرایط مذکور صدق کند و برای هر } x \text{ در دامنه آن، } g(x) \geq g(x).$$

او نشان داد که برآوردگر DEA سازگاری ضعیفی دارد و DEA برآوردگر بیشینه درست‌نمایی از مرز اصلی است.

بنکر [۶] اولین کسی بود که به ارزیابی استنباط‌های آماری روی مرز DEA پرداخت. تا قبل از ارزیابی نتایج بنکر، DEA به خاطر ناپارامتری بودن و عدم توانایی در ارزیابی تحلیل‌های آماری، مورد حمله و انتقاد بسیاری از محققان قرار گرفته بود. اگرچه بنکر نتایج قابل توجهی را برای DEA به دست آورد اما هیچ صحبتی از سرعت همگرایی مرز DEA به مرز ناشناخته اصلی نکرد.

کارستولو و همکارانش [۷ و ۸] ضمن بررسی مجدد کار بنکر، سازگاری و سرعت همگرایی مرز DEA را به مرز ناشناخته اصلی در حالت یک خروجی، بررسی کردند. ضمناً آن‌ها نشان دادند که تحت برخی شرایط عمومی ضعیف، مرز DEA برآوردگر درست‌نمایی بیشینه از مرز اصلی است و منطبق بر نتیجه‌ای بود که بنکر به دست آورده بود.

سرعت همگرایی یک برآوردگر از پارامتر اصلی به معیار اندازه‌گیری اختلاف آن‌ها بستگی دارد. کارستولو و همکارانش از روش معمول اندازه لبگ تفاضل متقارن برای سنجش این اختلاف استفاده کردند:

$$d_{\Delta}(\Psi, \hat{\Psi}_{\text{DEA-VRS}}) = \text{mes}(\Psi, \hat{\Psi}_{\text{DEA-VRS}}) \quad (1)$$

که

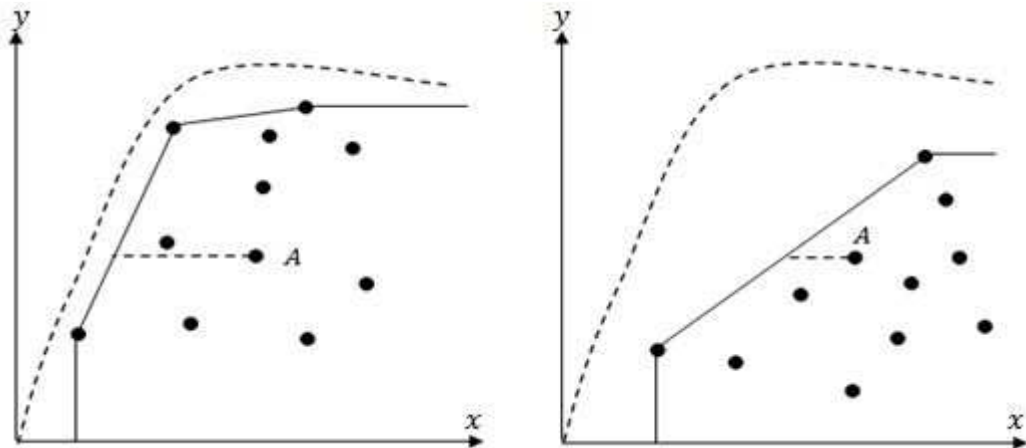
$$\Psi = \{(X, Y) \in R^{p+q} \mid X \text{ قابل تولید باشد} \mid Y\} \quad (2)$$

مجموعه امکان تولید اصلی و ناشناخته و

$$\Psi_{\text{DEA-VRS}} = \left\{ (X, Y) \mid X \geq \sum_{j=1}^n \eta_j X_j, Y \leq \sum_{j=1}^n \eta_j Y_j, \sum_{j=1}^n \eta_j = 1 \right\}, \quad (3)$$

مرز DEA تحت فرض بازده به مقیاس متغیر است.

آن‌ها نشان دادند که $E_p(n^{\frac{2}{p+2}} d_{\Delta}(\Psi, \hat{\Psi}_{DEA}))$ به طور مجانبی کراندار است؛ به این معنا که برای n های خیلی بزرگ، اختلاف بین Ψ و $\hat{\Psi}_{DEA}$ برابر با $O(n^{\frac{2}{p+2}})$ است. علاوه بر این، آن‌ها نشان دادند که هر برآورد دیگری با مجموعه امکان تولید محدب و مرزهای یکنوا نمی‌تواند با سرعت بهتری نسبت به مرز DEA، به مرز اصلی همگرا باشد. با وجود این خصوصیات بارز برای مرز DEA، باید متذکر شویم که سرعت همگرایی آن کند بوده؛ علاوه بر آن با افزایش تعداد ورودی‌ها این نرخ کمتر نیز می‌شود. به دو شکل زیر دقت کنید که در آن واحد A در دو نمونه ظاهر شده است.



شکل ۳. نمونه‌ی دوم

شکل ۲. نمونه‌ی اول

کارایی واحد A در نمونه‌ی دوم به طور محسوسی تغییر یافته است. این تغییر در مقدار کارایی ناشی از طبیعت ناپارامتری مدل DEA می‌باشد. مرز DEA وابسته به نمونه است و با تغییر نمونه مرز قبلی فرو می‌پاشد. اصطلاحاً گفته می‌شود که کارایی DEA نسبت به اندازه‌ی نمونه‌گیری، حساس است. البته همه ضعف این مدل، به دلیل ناپارامتری بودن آن نیست بلکه به اندازه نمونه نیز بستگی دارد و اینکه ما تنها یک نمونه از جامعه ناشناخته را در دست داریم. به هر حال یکی از نقاط ضعف اساسی مدل‌های DEA است که برخی محققان بدون در نظر گرفتن آن و بدون هیچ اقدامی در افزایش میزان دقت DEA به استفاده از این مدل‌ها برای تخمین کارایی دست می‌زنند [۹].

۳ ساختار کلی بوت استرپ

فعالیت یک سازمان تولیدی، به وسیله مجموعه امکان تولید Ψ که متشکل از تمام نقاط قابل تولید است؛ ارزیابی می‌شود.

$$\Psi = \{(X, Y) \in R^{p+q} \mid X \text{ قابل تولید باشد} \} \quad (۴)$$

این مجموعه اشتراک دو مجموعه امکان ورودی و مجموعه امکان خروجی است.

$$X(y) = \{x \in \mathbb{R}^p | (x, y) \in \mathbb{R}^{p+q}\} \quad (5)$$

$$Y(x) = \{y \in \mathbb{R}^q | (x, y) \in \mathbb{R}^{p+q}\} \quad (6)$$

تعاریف فوق منجر به مفهومی چون مرز کارای ورودی $\partial X(y)$ و مرز کارای خروجی $\partial Y(x)$ می شود که به ترتیب زیر مجموعه هایی از $X(y)$ و $Y(x)$ هستند.

$$\partial X(y) = \{x | x \in X(y), \theta x \in X(y); 0 < \theta < 1\} \quad (7)$$

$$\partial Y(x) = \{y | y \in Y(x), \beta x \in X(y); 1 < \beta\} \quad (8)$$

با توجه به تعریف فوق، مقادیر کارایی ورودی-محور و خروجی-محور در نقطه (x_k, y_k) به صورت زیر تعریف می شود.

$$\theta_k = \text{Min}\{\theta | \theta x_k \in X(y_k)\}, \quad (9)$$

$$\beta_k = \text{Max}\{\beta | \beta y_k \in Y(x_k)\}. \quad (10)$$

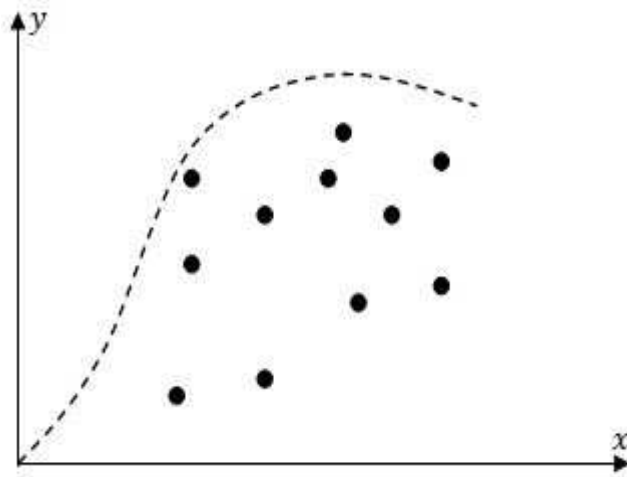
مقدار کارایی هر واحد، آن را به سطح کارای روی مرز منتقل می کند. این سطح کارایی ورودی-محور (خروجی-محور) موقعیت ورودی (خروجی) تصویر شده روی مرز است که با $X^\partial(x_k | y_k)$ نشان می دهیم:

$$X^\partial(x_k | y_k) = \theta x_k. \quad (11)$$

متأسفانه مجموعه امکان تولید مشخص نیست. بنابراین مجموعه امکان ورودی $X(y)$ ، مجموعه امکان خروجی $Y(x)$ ، مرز کارایی ورودی $\partial X(y)$ ، مرز کارایی خروجی $\partial Y(x)$ ، مقادیر کارایی ورودی-محور θ_k و خروجی-محور β_k ، قابل شناسایی و اندازه گیری نیستند.

تنها اطلاعاتی که از جامعه ناشناخته Ψ در اختیار داریم؛ یک نمونه از داده هاست. بهترین کار، مشخص کردن یک مرز به عنوان مرز کارایی و ارزیابی واحدها نسبت به آن است. واضح است مقادیر کارایی که از مقایسه با این مرز حاصل می شود؛ به میزان سازگاری و میزان اریبی این مرز از مرز اصلی بستگی دارد. هر قدر که این مرز برآورد درستی از مرز کارایی اصلی باشد؛ مقادیر کارایی حاصل به همان اندازه به مقادیر کارای اصلی نزدیک تر است.

ارایه یک برآورد درست از پارامتر اصلی نه تنها به ساختار اصلی برآورد گر بستگی دارد، بلکه به نمونه در دست بررسی نیز بستگی دارد. اینکه یک نمونه چه قدر می تواند اطلاعاتی را در اختیار برآورد گر قرار دهد؛ باید مورد بحث و بررسی قرار گیرد. چیزی که واضح است، تعداد کم نمونه مشاهده شده و به عبارت دیگر کوچک بودن حجم نمونه برای برآورد درست می باشد. بنابراین باید دنبال روشی برای تولید داده از جامعه ی ناشناخته باشیم تا مشکل کمبود نمونه را رفع کنیم. بوت استرپ که توسط افرون [۱] معرفی و توسط افرون تیشیرانی [۱۰] توسعه یافت؛ بر پایه شبیه سازی فرآیند تولید داده ها (DGP) بنا نهاده شده است.



شکل ۴. یک نمونه‌ی مشاهده شده از جامعه‌ی ناشناخته

فرض کنید $\chi = \{(x_i, y_i) | i = 1, \dots, n\}$ تنها نمونه‌ای از جامعه اصلی است که در اختیار داریم. با توجه به توضیحات بالا مطمئناً برآورد مرز کارایی با این داده‌ها نتیجه‌ی مطمئنی حاصل نخواهد شد. فرض کنید داده‌های نمونه به وسیله روش تولید داده ρ ، از جامعه اصلی استخراج شده باشد. اگر بتوانیم این فرآیند استخراج را شناسایی کنیم؛ می‌توانیم مشکل کمبود نمونه را با تولید داده‌های بیشتر حل نماییم. برای برآورد کردن پارامتر مشخصی از جامعه نیاز به یک روش مشخص $\mathcal{M}$ داریم تا به کمک آن بتوانیم این پارامتر را به دست آوریم. به کمک روش $\mathcal{M}$ برآوردی از مجموعه‌های $\hat{\Psi}$ ، $\widehat{X}(y)$ و $\widehat{\partial X}(y)$ را به دست می‌آوریم. بنابراین طبق (۹) کارایی ورودی-محور واحد k ام به صورت زیر قابل برآورد است.

$$\hat{\theta}_k = \text{Min} \left\{ \theta \mid \theta x_k \in \widehat{X}(y_k) \right\} \quad (12)$$

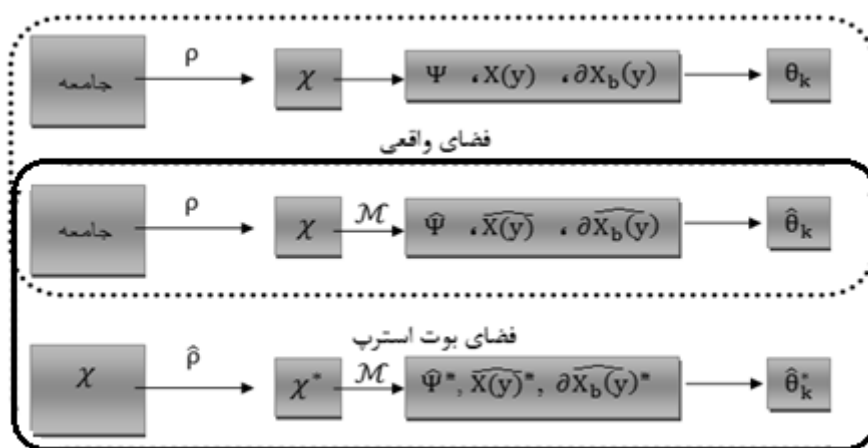
$\hat{\theta}_k$ برآوردی از مقدار واقعی θ_k است. مقدار درستی این برآورد به دو عامل روش $\mathcal{M}$ و حجم نمونه بستگی دارد. باید این نکته را یادآور شد که خصوصیات نمونه‌ای $\hat{\Psi}$ ، $\widehat{X}(y)$ و $\widehat{\partial X}(y)$ در نتیجه $\hat{\theta}_k$ ، به مولد داده ρ بستگی دارد که ناشناخته است. فرض کنیم با اطلاعاتی که از ρ داریم؛ با استفاده از داده‌های χ ، می‌توانیم برآوردگر معقول $\hat{\rho}$ از ρ را بسازیم. مجموعه داده‌های را که به وسیله $\hat{\rho}$ تولید شده است؛ در نظر می‌گیریم. این نمونه‌ی ساختگی به کمک روش $\mathcal{M}$ ، مجموعه‌های $\hat{\Psi}^*$ ، $\widehat{X}(y)^*$ و $\widehat{\partial X}(y)^*$ متناظر با مجموعه ساخته شده از χ^* را تعریف می‌کند. به ویژه اینکه برای واحد تحت بررسی (x_k, y_k) اندازه کارایی $\hat{\theta}_k^*$ ، به صورت زیر به دست می‌آید:

$$\hat{\theta}_k^* = \text{Min} \left\{ \theta \mid \theta x_k \in \widehat{\partial X}(y_k)^* \right\}. \quad (13)$$

از معلوم بودن $\hat{\rho}$ توزیع نمونه‌ای $\hat{\Psi}^*$ ، $\widehat{X(y)}^*$ و $\widehat{\partial X(y)}^*$ به طور مشروط روی χ مشخص است. هرچند شاید به طور تحلیلی، محاسبه آن سخت باشد. در مرحله بعد با به کار بستن روش مونت کارلو، توزیع نمونه‌ای به آسانی قابل تخمین است.

در اینجا نمونه‌ای اصلی χ در نقش جامعه قرار می‌گیرد و کارایی $\hat{\theta}_k^*$ یک برآورد از مقدار مشخص θ_k است. به عبارت دیگر در فضای بوت استرپ، χ به عنوان جامعه مشخص و χ^* یک نمونه از آن در فضای بوت استرپ است، اما $\hat{\theta}_k^*$ تنها از یک نمونه حاصل شده (شکل ۵). برای برآورد بهتر، دوباره اقدام به ساخت نمونه‌های ساختگی می‌کنیم. با استفاده از $\hat{\rho}$ ، B نمونه χ_b^* ، $b=1, \dots, B$ را تولید می‌نماییم. روش $\mathcal{M}$ را برای هر کدام از این نمونه‌های ساختگی به کار می‌گیریم تا هر کدام از این نمونه‌های ساختگی، برآوردهای $\hat{\Psi}_b^*$ ، $\widehat{X_b(y)}^*$ و $\widehat{\partial X_b(y)}^*$ را مشخص کند. در مرحله‌ی بعد با استفاده از برآوردگر اصلی، برای هر واحد تحت بررسی (x_k, y_k) ، مقادیر متناظر کارایی $\{\hat{\theta}_{k,b}^*\}_{b=1}^B$ را تعیین می‌کنیم. تابع چگالی تجربی مربوط به $\{\hat{\theta}_{k,b}^*\}_{b=1}^B$ ، تقریب مونت کارلو از توزیع θ_k^* ، روی $\hat{\rho}$ است. اگر $\hat{\rho}$ یک برآوردگر منطقی و سازگار از ρ باشد؛ آن‌گاه توزیع بوت استرپ، توزیع اصلی نامشخص را برآورد می‌کند (شبیه‌سازی می‌کند). به عبارت دیگر اگر $\hat{\rho}$ برآورد درستی از ρ باشد آن‌گاه رفتار $\hat{\theta}_k^*$ نسبت به $\hat{\theta}_k$ در فضای بوت استرپ، شبیه به رفتار $\hat{\theta}_k$ نسبت به θ_k در فضای واقعی خواهد بود. با در نظر گرفتن جزئیات بیشتر می‌توان گفت، برای اندازه کارایی θ_k متناظر با واحد تحت بررسی (x_k, y_k) داریم:

$$(\hat{\theta}_k^* - \hat{\theta}_k) | \hat{\rho} \sim (\hat{\theta}_k - \theta_k) | \rho \quad (14)$$



شکل ۵. فرآیند تولید داده

که θ_k ، $\hat{\theta}_k$ و $\hat{\theta}_k^*$ از روابط (۹)، (۱۲) و (۱۳) حاصل می‌شود. به یاد داشته باشید که رابطه‌ی (۱۴) زمانی برقرار است که $\hat{\rho}$ برآورد سازگاری از ρ باشد.

محاسبه‌ی میزان اریبی یک برآورد گر، می‌تواند مقدار برآورد شده را بهبود بخشد و برآورد بهتری از پارامتر مورد نظر ارائه کند. پس

$$\text{bias}_{\rho,k} = E(\hat{\theta}_k) - \theta_k \quad (15)$$

θ_k نامشخص است و از این رو مقدار اریبی نیز مجهول است. در فضای بوت استرپ مقدار اریبی برآورد گر $\hat{\theta}_k^*$ مشخص است و می‌توان آن را اندازه گیری نمود.

$$\text{bias}_{\hat{\rho},k} = E(\hat{\theta}_k^*) - \hat{\theta}_k \quad (16)$$

مقدار مورد انتظار برای $\hat{\theta}_k^*$ را می‌توان با تعریف مونت-کارلو از آن جایگزین کرد.

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

بنابراین

$$\text{bias}_{\hat{\rho},k} = \frac{1}{B} \sum_{b=1}^B \hat{\theta}_{k,b}^* - \hat{\theta}_k \quad (18)$$

با توجه به رابطه‌ی (۱۴) می‌توان برآوردی از (۱۵) را با استفاده از (۱۸) به دست آورد. بنابراین

$$\text{bias}_{\rho,k} \approx \frac{1}{B} \sum_{b=1}^B \hat{\theta}_{k,b}^* - \hat{\theta}_k \quad (19)$$

با اصلاح مقدار اریبی از برآورد گر اصلی ($\hat{\theta}_k$)، برآورد گر اریب-اصلاح شده‌ی بهتری حاصل می‌شود.

$$\begin{aligned} \tilde{\theta}_k &= \hat{\theta}_k - \text{bias}_{\rho,k} \approx \hat{\theta}_k - \left(\frac{1}{B} \sum_{b=1}^B \hat{\theta}_{k,b}^* - \hat{\theta}_k \right) \\ &\approx \hat{\theta}_k - \frac{1}{B} \sum_{b=1}^B \hat{\theta}_{k,b}^* \\ &\approx \hat{\theta}_k - \bar{\theta}_k^* \end{aligned} \quad (20)$$

عبارت برآورد گر اریب-اصلاح شده را از این جهت به کار می‌گیریم که مقدار اریبی به دست آمده برای این برآورد گر مقدار دقیقی آن نیست؛ بلکه یک مقدار تقریبی از آن است. از این رو مقدار اریبی برآورد گر از بین نمی‌رود و تنها اصلاح می‌شود.

خطای استاندارد $\hat{\theta}_k^*$ برابر است با

$$\widehat{\text{se}}_{\rho,k} = \left\{ \frac{1}{B-1} \sum_{b=1}^B (\hat{\theta}_{k,b}^* - \bar{\theta}_k^*)^2 \right\}^{\frac{1}{2}}. \quad (21)$$

سرانجام بعد از اعمال مقدار اریبی به تابع توزیع به دست آمده، بازه‌های اطمینان را محاسبه می‌کنیم. در نظر داریم که تابع توزیع تخمین زده شده طوری قرار گیرد که مرکز آن $\tilde{\theta}_k$ باشد. بنابراین باید تابع توزیعی را که برای $\hat{\theta}_k$

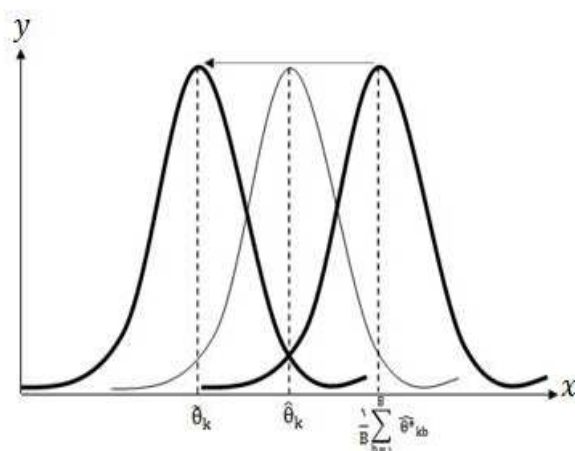
به دست می آید؛ به اندازه $\widehat{bias}_k$ به سمت چپ منتقل کنیم (اصلاح مقدار اریبی $\hat{\theta}_k$). تابع توزیع به دست آمده به مرکز تقریبی $\frac{1}{B} \sum_{b=1}^B \hat{\theta}_{kb}^*$ می باشد و برای مرکزیت قرار دادن $\hat{\theta}_k$ ، باید مقدار اریبی را اصلاح کرد. بنابراین بعد از به دست آوردن تابع توزیع $\{\hat{\theta}_{kb}^*\}_{b=1}^B$ ، آن را به اندازه $\widehat{bias}_k$ ، به سمت چپ منتقل کنیم (شکل ۶). برآوردگر اصلاح شده به صورت زیر خواهد بود.

$$\tilde{\theta}_{k,b}^* = \hat{\theta}_{k,b}^* - \widehat{bias}_k. \quad (22)$$

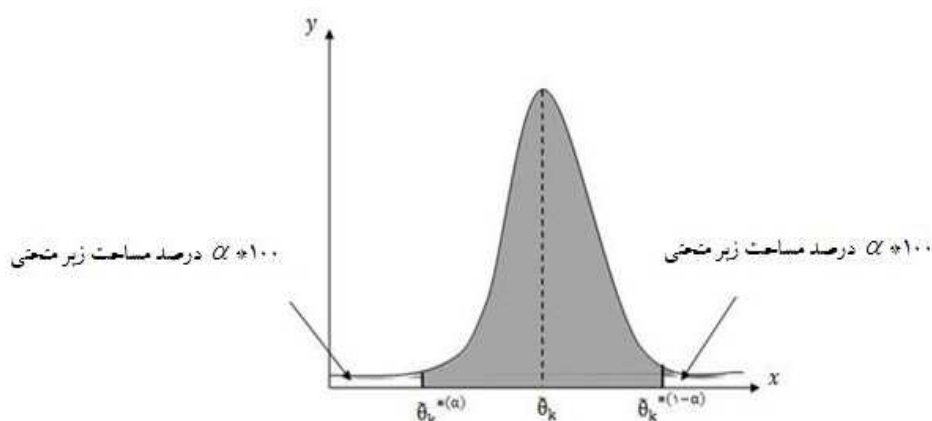
بازه ی اطمینان درصدی معمولی برای $\hat{\theta}_k$ که در آن طول بازه $1-\alpha$ است؛ به صورت زیر تعریف می شود:

که $\tilde{\theta}_k^{*(\frac{\alpha}{2})}$ ، $100 - \left(\frac{\alpha}{2}\right)$ ام تابع چگالی تجربی $\{\tilde{\theta}_{k,b}^*\}_{b=1}^B$ می باشد (برای تفسیر شهودی آن به شکل فرضی ۷ مراجعه کنید).

اگر تابع چگالی تجربی $\hat{\theta}_{kb}^*$ اریب باشد؛ شاید بهتر باشد که $\tilde{\theta}_k$ در مرکز تابع توزیع باشد [۴].

$$(\hat{\theta}_{k,low}, \hat{\theta}_{k,up}) = \left(\tilde{\theta}_k^{*(\frac{\alpha}{2})}, \tilde{\theta}_k^{*(1-\frac{\alpha}{2})} \right). \quad (23)$$


شکل ۶. اصلاح اریبی تابع توزیع



شکل ۷. ساخت بازه ی اطمینان

اصلاح مقدار اریبی، ممکن است موجب بروز خطای میانگین ریشه دوم بزرگی در برآورد و هم چنین باعث افزایش واریانس برآورد گر شود. اگر مقدار اریبی محاسبه شده (۱۹) در مقایسه با خطای استاندارد (۲۱) مقدار کوچکی باشد. در این صورت نیاز به اصلاح مقدار اریبی نیست و بهتر است توزیع به مرکزیت $\hat{\theta}_k$ باشد [۱۱]. افرون و تیشرانی [۱۱] پیشنهاد دادند که اگر نسبت $\frac{\widehat{bias}_{p,k}}{\widehat{se}_{p,k}}$ از ۰/۲۵ تجاوز نکند؛ آن گاه بهتر است اصلاح مقدار اریبی صورت نگیرد و $\hat{\theta}_k$ مرکز توزیع تجربی قرار گیرد. تنها سوالی که در اینجا باقی می ماند روش تولید داده $\hat{\rho}$ است.

۴ تخمین یک مولد منطقی از مولد داده ی واقعی

در بخش قبل به معرفی کلیت کار بوت استرپ پرداختیم. تنها بحث باقی مانده از بخش قبل، تخمین یک برآورد منطقی از مولد داده ρ است. بوت استرپ در چارچوب مدل های مرزی ناپارامتری توسط سیمار [۱۲] معرفی و توسط سیمار - ویلسون [۴] توسعه یافت. سیمار و ویلسون [۴ و ۱۳]، یک چارچوب کلی برای روش تولید داده در مدل ها مرزی ناپارامتری ترسیم کردند.

مجموعه داده ی $\chi_n = \{(x_i, y_i) | i = 1, \dots, n\}$ به وسیله مولد داده ی ρ از جامعه ناشناخته ی Ψ استخراج شده است. این مولد داده ی ρ به صورت کامل به وسیله ی Ψ و تابع چگالی احتمال $f(x, y)$ مشخص می شود.

$$\rho = \rho(\Psi, f(x, y)) \quad (24)$$

اگر $\hat{\rho}(\chi_n)$ برآورد درستی از مولد داده اصلی باشد؛ آن گاه

$$\hat{\rho}(\chi_n) = \rho(\hat{\Psi}, \hat{f}(x, y)). \quad (25)$$

در فضای واقعی $f(x, y)$ و Ψ ناشناخته است. تنها نمونه χ مشخص است و باید برای برآورد ρ ، Ψ و θ استفاده شود.

یک فضای مجازی شبیه سازی شده مانند فضای بوت استرپ را در نظر بگیرید. آنچه در فضای بوت استرپ اتفاق می افتد؛ اتفاقات فضای واقعی را تقلید می کند. در فضای بوت استرپ، $\hat{\rho}$ ، $\hat{\Psi}$ و $\hat{\theta}_k$ مشخص هستند و جایگزین ρ ، Ψ و θ_k در فضای واقعی می شوند. در فضای واقعی، ρ مولد داده ی واقعی و مجهول و $\hat{\rho}$ مولد داده ی برآورد شده از ρ ، اما در فضای بوت استرپ، $\hat{\rho}$ مولد داده ی واقعی و مشخص می باشد.

در اینجا روش $\mathcal{M}$ که در بخش قبل برای برآورد Ψ و سایر مجموعه های متناظر با آن استفاده شد؛ DEA است. مجموعه های برآورد شده توسط روش $\mathcal{M}$ که در بخش قبل دیدیم را بار دیگر و این بار به کمک DEA بیان می کنیم.

$$\hat{\Psi}_{\text{DEA-VRS}} = \left\{ (X, Y) \mid X \geq \sum_{j=1}^n \eta_j X_j, Y \leq \sum_{j=1}^n \eta_j Y_j, \sum_{j=1}^n \eta_j = 1 \right\} \quad (26)$$

$$\begin{cases} \widehat{X(y_k)} = \{x \in \mathbb{R}_+^p | (x, y_k) \in \hat{\Psi}_{DEA}\} \\ \widehat{Y(x_k)} = \{y \in \mathbb{R}_+^q | (x, y_k) \in \hat{\Psi}_{DEA}\} \end{cases} \quad (27)$$

$$\partial \widehat{X(y_k)} = \{x | x \in X(y_k), \theta x \notin X(y) \quad \forall 0 < \theta < 1\} \quad (28)$$

$$\hat{\theta}_k = \text{Min}\{\theta | \theta x_k \in X(y_k)\} \quad (29)$$

و یا

$$\hat{\theta}_k = \text{Min} \quad \theta,$$

s.t.

$$\begin{aligned} Y_k &\leq \sum_{j=1}^n \gamma_j Y_j, \\ \theta X_k &\geq \sum_{j=1}^n \gamma_j X_j, \\ \sum_{j=1}^n \gamma_j &= 1, \\ \gamma_j &\geq 0, \quad i = 1, \dots, n. \end{aligned} \quad (30)$$

و

$$\hat{X}^\partial(x_k | y_k) = \hat{\theta}_k x_k \quad (31)$$

۴-۱ بوت استرپ ساده و تشریح عدم استفاده از این روش برای مدل‌های مرزی

در این بخش ابتدا به تشریح فرآیند بوت لسترپ ساده (خام) می‌پردازیم و نشان می‌دهیم استفاده از این روش که در مقاله سعید عبادی از آن استفاده شده منجر به نتایج نامعتبری خواهد شد.

۴-۱-۱ بوت استرپ ساده

از آنجا که ρ ، مجموعه داده‌های $\{(X_i, Y_i) | i = 1, \dots, n\}$ را تولید می‌کند؛ بهترین برآوردگر تابع توزیع اصلی در این زمینه، تابع توزیع تجربی است که در هر نقطه جرم احتمال $\frac{1}{n}$ را اختیار می‌کند. در فضای یک بعدی، تابع توزیع تجمعی تجربی، به صورت زیر محاسبه می‌شود.

$$F_n(x) = \sum_{i=1}^n I(x_i \leq x) \quad (32)$$

در این صورت، نمونه ساختگی را می‌توان با انتخاب تصادفی و با جای‌گذاری از مجموعه χ_n به دست آورد. ایده اصلی نمونه‌سازی از طریق انتخاب تصادفی و با جای‌گذاری، به این حقیقت مربوط است که در صورت بزرگ بودن حجم نمونه، توزیع تجربی $F_n(x)$ به توزیع اصلی ناشناخته $F(x)$ نزدیک می‌شود.

با توجه به تعریف (۹) از کارایی، کارایی نقطه (X_k, Y_k) که طی این فرآیند ثابت در نظر گرفته می شود؛ (فرآیند بوت استرپ برای هر واحد باید به طور جداگانه عملی شود). با توجه به نمونه استخراج شده به صورت زیر محاسبه می شود.

$$\begin{aligned} \hat{\theta}_k^* &= \text{Min } \theta, \\ \text{s.t.} \\ Y_k &\leq \sum_{j=1}^n \gamma_j Y_j^*, \\ \theta X_k &\geq \sum_{j=1}^n \gamma_j X_j^*, \\ \sum_{j=1}^n \gamma_j &= 1, \\ \gamma_j &\geq 0, \quad i = 1, \dots, n. \end{aligned} \quad (33)$$

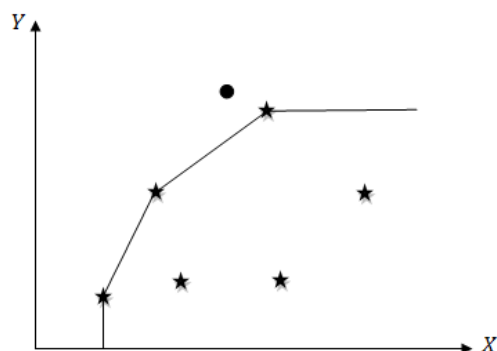
تابع توزیع اصلی ناشناخته است و چون تابع توزیع تجربی هر نمونه از جامعه به تابع توزیع جامعه همگراست؛ بنابراین تابع تجربی، برآورد درستی از تابع اصلی F و نیز $\hat{\Psi}$ برآورد سازگاری از Ψ است (مجموعه $\hat{\Psi}$ به وسیله DEA برآورد شده که در دیدیم، برآورد درستی از مجموعه امکان تولید است). بنابراین مولد داده $(\hat{\rho}(\chi_n), \hat{\rho}(\chi_n))$ یک برآورد درست از مولد داده اصلی ρ است [۴].

θ_k در فضای واقعی نامشخص و $\hat{\theta}_k$ مقداری مشخص و به عنوان تخمینی از آن در نظر گرفته می شود. در فضای واقعی نیز $\hat{\theta}_k$ مقدار درست کارایی و مشخص و $\hat{\theta}_k^*$ تخمینی از $\hat{\theta}_k$ در فضای بوت استرپ است. $\hat{\rho}(\chi_n)$ یک برآورد درست از ρ است. پس رابطه (۱۴) برقرار می باشد؛ یعنی در فضای بوت استرپ رفتار $\hat{\theta}_k^*$ نسبت به $\hat{\theta}_k$ ، مشابه رفتار $\hat{\theta}_k$ نسبت به θ_k در فضای واقعی است. از این رو عملیات توضیح داده شده در بخش قبل، برای یافتن تخمین بهتر کارایی و ساختن بازه های اطمینان، عیناً انجام می گیرد.

استراتژی معرفی شده بسیار جالب است اما متأسفانه برای مدل های مرزی کاربرد ندارد؛ چرا که اندازه ناکارایی فارل را که همان انحراف شعاعی x از $\partial X(y)$ است؛ به درستی منعکس نمی کند [۴].

برای برخی مجموعه های تولید شده χ_n^* از توزیع تجربی، $\partial X(y)$ قابل تعریف نیست. بنابراین مقدار متناظر با آن قابل اندازه گیری نخواهد بود.

شکل ۱ نشان می دهد که برای واحد تحت بررسی، اندازه ی کارایی ورودی محور قابل محاسبه نیست. وقتی در چارچوب رگرسیون کار می کنیم؛ بهتر است؛ بوت استرپ روی باقی مانده ها عملی شود [۱۴ و ۱۵] که در اینجا باقیمانده ها، همان مقادیر θ_k می باشند.



شکل ۸. شکست مدل در محاسبه‌ی کارایی ورودی محور

■: واحد تحت بررسی

*: نقاط تولید شده به وسیله‌ی انتخاب تصادفی و با جای گذاری از نمونه اصلی

DGP را می‌توان به این صورت تعریف کرد که برای یک مقدار مشخص y ، می‌دانیم $x \in X(y)$ به دلیل ناکارا بودن این نقطه، نمی‌تواند متعلق به $\partial X(y)$ باشد و روی یک شعاع دور شونده از $X^\partial(x_k|y_k)$ قرار دارد. بنابراین یک واحد مشخص (x_k, y_k) ، به صورت نقطه‌ای تولید شده در نظر گرفته می‌شود که سطح خروجی آن y_k و ورودی آن به یک نسبت تصادفی، افزایش یافته است؛ یعنی $x_i^* = \frac{X^\partial(x_i|y_i)}{\theta_i}$. فرض کنید که فرآیند تولید داده‌ها، مقادیر ناکارایی‌های $\theta_k \in (0, 1]$ زیر را تولید کرده باشد:

$$(\theta_1, \dots, \theta_n) \sim \text{i.i.d.} F \quad (34)$$

که F تابع چگالی روی $(0, 1]$ می‌باشد. همان‌طور که در بالا گفته شد هر نقطه را در سطح ثابت خروجی y_i به صورت نقطه‌ی $x_i^* = \frac{\hat{X}^\partial(x_i|y_i)}{\theta_i^*}$ تولید شده در نظر می‌گیریم که θ_i اندازه‌ی شعاعی فاصله آن از سطح کارایی $X^\partial(x_i|y_i)$ باشد. $\mathcal{P}_i$ را مولد داده‌ی متناظر با (x_i, y_i) در نظر می‌گیریم. بنابراین هر نقطه را می‌توان با $X^\partial(x_k|y_k)$ و F نمایش داد:

$$\rho_i = (X^\partial(x_i|y_i), F) \quad (35)$$

که در نهایت، کل مولد داده، برابر با مجموعه $\rho = (\rho_1, \dots, \rho_n)$ خواهد بود. اگر $\partial X(y_k)$ و $X^\partial(x_k|y_k)$ مشخص باشد؛ مقدار $\theta_k = \frac{X^\partial(x_k|y_k)}{x_k}$ نیز مشخص می‌شود و F را با تابع توزیع تجربی آن‌ها جای‌گزین می‌کنیم (به عنوان برآوردی از آن). متأسفانه $\partial X(y_k)$ نامشخص است اما می‌توانیم از $\widehat{\partial X(y_k)}$ و $\hat{X}^\partial(x_k|y_k)$ به عنوان تخمینی از $\partial X(y_k)$ و $X^\partial(x_k|y_k)$ استفاده نماییم و مقدار $\hat{\theta}_k$ را محاسبه کنیم. ساده‌ترین برآورد از F ، توزیع تجربی $\hat{\theta}_k$ ‌ها می‌باشد:

$$\hat{F}(t) = \begin{cases} n^{-1} & \text{if } t = \hat{\theta}_k, i = 1, \dots, n \\ 0 & \text{O.W} \end{cases} \quad (36)$$

بعد از مشخص شدن $\hat{F}$ و $\hat{X}^\partial(x_i|y_i)$ تعریف می‌کنیم:

$$\hat{\rho}_i = (\hat{X}^\partial(x_i|y_i), \hat{F}). \quad (37)$$

با استفاده از تعریف فوق، می توان با سطح خروجی y_k ، x_k^* را ساخت. حال به طور تصادفی و با جای گذاری مقادیر θ_i^* را از θ_i انتخاب کنیم:

$$(\theta_1^*, \dots, \theta_n^*) \sim \text{i.i.d.F.} \quad (38)$$

بنابراین ورودی های تولید شده به صورت زیر خواهد بود

$$x_i^* = \frac{\hat{X}^\partial(x_i|y_i)}{\theta_i^*} = \frac{\hat{\theta}_i}{\theta_i^*} x_i. \quad (39)$$

ادامه کار همان روند توضیح داده شده در بخش قبل است یعنی به کارگیری روش DEA و به دست آوردن کارایی $\hat{\theta}_i^*$ و ادامه همین روند به تعداد B بار و سپس محاسبه ی مقدار اریبی، اصلاح مقدار اریبی و ساختن بازه های اطمینان.

فرآیند تولید داده بالا، بازهم برای مدل های با مرز ناپارامتری قابل استفاده نیست. سیمار - ویلسون [4] دلیل ناکارآمدی این روش را برشمرده اند که در اینجا به آن اشاره می کنیم.

۴-۱-۲ اشکالات موجود در روش بوت استرپ ساده

برآوردگر DEA و FDH، تعداد زیادی واحد با کارایی واحد تولید می کنند که تعداد این واحدها با افزایش تعداد ورودی ها، افزایش می یابد. در نتیجه F_i برآورد کم ارزشی از F را در نزدیکی کران بالا ($\theta = 1$) ارائه می کند. افرون و تیشیرانی [۱۰] نشان دادند که در نزدیکی کران بالا، تابع توزیع تجربی، برآورد سازگاری از F نمی دهد. اگر F روی $(0, 1)$ به طور پیوسته فرض شود؛ در نزدیکی $\theta = 1$ ، تجمع زیادی ایجاد خواهد شد که این مساله باعث ناکارآمدی برآوردگر می شود. علاوه براین، برآورد F با $\hat{F}$ ، در قسمت خطوط کرانه ای وقتی که تکیه گاه کراندار باشد؛ کار مشکلی است. در تخمین مرز کارایی تنها کران های بالایی ($\theta = 1$) مشکل ساز است و اگر مشکل آن برطرف نشود برآوردگر بوت استرپ، ناسازگار خواهد بود [4].

سیمار - ویلسون [4] نشان دادند که در صورت استفاده از بوت استرپ ساده، خواهیم داشت:

$$\text{Prob}(\hat{\theta}^*(x, y) = \theta(x, y) | \chi_n) = 1 - (1 - n^{-1})^n > 0 \quad (40)$$

که

$$\lim_{n \rightarrow \infty} \text{Prob}(\hat{\theta}^*(x, y) = \theta(x, y) | \chi_n) = 1 - e^{-1} = 0.632 \quad (41)$$

مطلقاً هیچ دلیلی موجود نیست که مقدار احتمال بالا برابر 0.632 باشد؛ حتی بدون در نظر گرفتن هیچ مشخصه ای از فرآیند تولید داده، این امکان وجود ندارد. اگر تابع چگالی احتمال مقادیر کارایی، پیوسته باشد؛ مقدار احتمال بالا برابر صفر خواهد بود [4].

$$\text{Prob}(\hat{\theta}^*(x, y) = \theta(x, y)) = 0 \quad (42)$$

هم چنین اگر تابع چگالی احتمال اصلی در هر نقطه دارای مقدار جرم $\frac{1}{n}$ باشد؛ هیچ دلیلی وجود ندارد که این مقدار برابر $0/632$ باشد.

برای رفع این مشکل، سیمار و ویلسون یک نسخه‌ی هموار از بوت استرپ را ارایه کردند که به تفصیل در بخش ۴-۲ به آن خواهیم پرداخت.

۴-۲ بوت استرپ هموار

در بخش قبل در مورد بوت استرپ ساده بحث کردیم. دیدیم که این روش برای تخمین مرز کارایی ناکارآمد است و منجر به برآورد ناسازگاری از مرز کارایی می‌شود. مشکل به وجود آمده ناشی از این حقیقت می‌باشد که مرز کارایی واقعی پیوسته را با یک تخمین ناپیوسته از آن که در هر نقطه احتمال $\frac{1}{n}$ را اختیار می‌کند؛ جای‌گزین کرده‌ایم.

سیمار و ویلسون [۴] یک نسخه هموار از بوت استرپ را ارایه کردند که تفاوت آن با نسخه قبلی در فرآیند تولیدهاست. در این روش یک برآوردگر هموار از تابع چگالی احتمال مقادیر کارایی ساخته می‌شود. که تولید داده به کمک آن صورت می‌گیرد.

فرض کنیم تابع چگالی احتمال مقادیر کارایی $f(\theta)$ پیوسته باشد. برآوردگر چگالی هسته‌ی استاندارد از $f(\theta)$ به صورت زیر تعریف می‌شود.

$$\hat{f}_h(t) = (nh)^{-1} \sum_{i=1}^n K\left(\frac{t - \hat{\theta}_i}{h}\right) \quad (43)$$

که h پارامتر هموارسازی یا پهنای باند و $K(\circ)$ تابع هسته است؛ به طوری که $K(t) = K(-t)$ و $\int_{-\infty}^{\infty} K(t) dt = 1$

و $\int_{-\infty}^{\infty} tK(t) dt = 0$. $\hat{f}_h(\theta)$ میانگین n تابع چگالی $K(\circ)$ به مرکزیت مقادیر کارایی $\hat{\theta}_i$ ها می‌باشد. برای مشخص

کردن برآوردی از تابع چگالی، دو انتخاب باید صورت بگیرد: اولی انتخاب تابع هسته و دومی انتخاب پارامتر هموارسازی. برای تابع هسته معمولاً تابع چگالی احتمال استاندارد (گوسی) انتخاب می‌شود (از اینجا به بعد تابع هسته‌ای که برای مولد داده استفاده می‌کنیم؛ تابع چگالی احتمال استاندارد (گوسی) است). دقت در انتخاب پارامتر هموارسازی از اهمیت بیشتری نسبت به انتخاب تابع چگالی هسته برخوردار است؛ چراکه تاثیر زیادی در کیفیت برآوردگر دارد [۹]. در ادامه در مورد چگونگی انتخاب پارامتر هموارسازی بحث خواهد شد.

با وجود ارایه‌ی یک برآورد هموار از تابع چگالی احتمال، بازهم شرایط مرزی $t < 1$ یکی از مشکلاتی است که به دلیل کراندار بودن فضا هم چنان پابرجاست [۴].

روش بازتابی تشریح شده به وسیله سیلورمن [۱۴]، ابزار ساده‌ای برای غلبه بر این مشکل است. برای هر نقطه $\hat{\theta}_i \leq 1$ ، تصویر متقارن آن نسبت به ۱، یعنی $\hat{\theta}_i - 1 \leq 1$ ، $i = 1, \dots, n$ در نظر گرفته می‌شود. سپس تابع چگالی هسته را از مجموعه $2n$ نقطه‌ای به صورت زیر برآورد می‌کنیم:

بهاری و بهکاران، استفاده از فرآیند شبیه سازی بوت استرپ برای برآورد مرز تولید کارای نامار استری بررسی مشکلات موجود در فرآیند آیدار شده در مقاله سیدعبادی

$$\hat{G}_h(t) = \frac{1}{rnh} \sum_{i=1}^n \left(K\left(\frac{t-\hat{\theta}_i}{h}\right) + K\left(\frac{r-t+\hat{\theta}_i}{h}\right) \right). \quad (44)$$

حال تعریف می کنیم

$$\hat{f}_{s,h}(t) = \begin{cases} r\hat{G}_h(t) & \text{اگر } t \leq 1 \\ 0 & \text{در غیر این صورت} \end{cases} \quad (45)$$

$\hat{f}_{s,h}(t)$ یک تابع متقارن حول ۱ است که سمت چپ آن منطبق بر تابع $f(\theta)$ می باشد. شوستر [۱۵] نشان داد برای هر $t \leq 1$ ، $\hat{f}_{s,h}(t)$ یک برآوردگر سازگار از f می باشد. حال با استفاده از $\hat{f}_{s,h}(t)$ اقدام به تولید نمونه های ساختگی می کنیم. فرض کنید، $\beta_1^*, \beta_r^*, \dots, \beta_n^*$ ، نمونه به دست آمده از $\hat{\theta}_1, \hat{\theta}_r, \dots, \hat{\theta}_n$ به وسیله انتخاب تصادفی با جای گذاری باشد. با اطلاعات مقدماتی آمار و احتمال می توان نشان داد:

$$t_i = \beta_i^* + h\epsilon_i^* \sim \hat{G}_{1,h}(t) = \frac{1}{n} \sum_{i=1}^n \frac{1}{h} K\left(\frac{t-\hat{\theta}_i}{h}\right) \quad (46)$$

که ϵ_i^* به صورت تصادفی و از توزیع نرمال استاندارد استخراج شده است. به طور مشابه برای نقاط منعکس شده می توان نشان داد:

$$t_i^R = r - \beta_i^* - h\epsilon_i^* \sim \hat{G}_{r,h}(t) = \frac{1}{n} \sum_{i=1}^n \frac{1}{h} K\left(\frac{t-r+\hat{\theta}_i}{h}\right). \quad (47)$$

حال با توجه به رابطه

$$\hat{f}_{s,h}(t) = \hat{G}_{r,h}(t) + \hat{G}_{1,h}(t) \quad (48)$$

و با در نظر گرفتن مولد تصادفی زیر

$$\tilde{\theta}_i^* = \begin{cases} \beta_i^* + h\epsilon_i^* & \beta_i^* + h\epsilon_i^* \leq 1 \\ r - \beta_i^* - h\epsilon_i^* & \text{O.W} \end{cases} \quad (49)$$

خواهیم داشت،

$$\tilde{\theta}_i^* \sim \hat{F}_{s,h}(t). \quad (50)$$

که به آسانی و با توجه به روابط اخیر قابل اثبات است. بنابراین نمونه جدید نیز از توزیع نمونه اصلی پیروی می کند.

سوالی که در اینجا پیش می آید این است که آیا روند تولید داده ی فوق یک سیر منطقی دارد؟ آیا روابط (۴۴) و (۴۷) باید برقرار باشند؟ و اینکه چرا باید این روابط برقرار باشند؟ چرا در قسمت قبل که تابع توزیع تجربی را به عنوان برآوردی از تابع توزیع اصلی قرار دادیم؛ به صورت تصادفی و با جای گذاری اقدام به تولید داده ها کردیم؟

نشان دادن این که رابطه‌ی (۵۰) برقرار است از ملزومات ادامه روند شبیه‌سازی است؛ چراکه گفتیم $\hat{f}_h(t)$ (نیمه سمت چپ تابع $\hat{f}_{s,h}(t)$) را به عنوان برآوردی از تابع $f(\theta)$ در نظر می‌گیریم. یعنی توزیع مقادیر کارایی بدین صورت است. بنابراین داده‌های شبیه‌سازی شده باید از این توزیع پیروی کنند. هم‌چنین، در بوت استرپ ساده، تابع توزیع تجربی که در آن هر نقطه احتمال $\frac{1}{n}$ را اختیار می‌کند؛ به عنوان برآوردی از تابع ناشناخته اصلی در نظر گرفته می‌شود و با انتخاب تصادفی و با جای‌گذاری احتمال انتخاب هر نقطه $\frac{1}{n}$ می‌باشد. بنابراین نمونه به دست آمده، مجدداً از توزیع تجربی پیروی می‌کند.

با بررسی بیشتر در خصوصیات آماری نمونه تولید شده به روابط زیر می‌رسیم [۴]

$$E(\tilde{\theta}_i^* | \hat{\theta}_1, \dots, \hat{\theta}_n) = \hat{\mu} \quad (51)$$

$$V(\tilde{\theta}_i^* | \hat{\theta}_1, \dots, \hat{\theta}_n) = \hat{\sigma}_{\theta}^2 + h^2 \quad (52)$$

که $\hat{\sigma}_{\theta}^2$ واریانس جامعه $\theta_1, \theta_2, \dots, \theta_n$ ، یعنی

$$\hat{\sigma}_{\theta}^2 = \frac{1}{n} \sum_{i=1}^n \left(\hat{\theta}_i^2 - \hat{\theta}_i \right) \quad (53)$$

و $\hat{\mu}$ میانگین نمونه می‌باشد. اگر برآورد گر هسته را مورد استفاده قرار دهیم واریانس دنباله بوت استرپ به صورت زیر اصلاح می‌شود.

$$\theta_i^* = \beta_i^* + \frac{1}{\left(1 + \frac{h^2}{\hat{\sigma}_{\theta}^2}\right)^{\frac{1}{2}}} (\tilde{\theta}_i^* - \bar{\beta}_1^*) \quad (54)$$

$$\bar{\beta}_1^* = \frac{1}{n} \sum_{i=1}^n \beta_i^* \quad \text{که}$$

می‌توان به وسیله محاسبات دستی مستقیمی نشان داد که

$$E(\theta_i^* | \hat{\theta}_1, \dots, \hat{\theta}_n) = \hat{\mu} \quad (55)$$

$$V(\theta_i^* | \hat{\theta}_1, \dots, \hat{\theta}_n) = \hat{\sigma}_{\theta}^2 \left(1 + \frac{h^2}{n(\hat{\sigma}_{\theta}^2 + h^2)}\right) \quad (56)$$

از این نظر که واریانس نمونه مجانباً کراندار است؛ دنباله θ_i^* به دست آمد از بوت استرپ هموار، خصوصیات بهتری نسبت به $\tilde{\theta}_i^*$ دارد [۴].

برای برآورد گر DEA، الگوریتم کامل بوت استرپ در گام‌های زیر خلاصه شده است (شکل ۹).

برای هر (x_k, y_k) ، $k=1, 2, \dots, n$ ، $\hat{\theta}_1, \hat{\theta}_2, \dots, \hat{\theta}_n$ را به وسیله برنامه‌ریزی خطی (۳۳) حل کنید.

با استفاده از روند توضیح داده شده در این بخش، برای بوت استرپ یک نمونه تصادفی به حجم n از مجموعه اصلی یعنی $\theta_1^*, \dots, \theta_n^*$ را تولید کنید.

Downloaded from jamlu.liau.ac.ir on 2025-11-05]

برای هر $\hat{\theta}_k$ که $k = 1, 2, \dots, n$ ، برآورد بوت استرپ $\hat{\theta}_{k,b}^*$ را با کمک حل برنامه‌ریزی خطی زیر به دست آورید.

مراحل ۲-۴ را به تعداد B بار تکرار کنید تا برای مجموعه برآورد زیر به دست آید:

برای مجموعه داده‌های بزرگ، انتخاب B محدود به منابع رایانه‌ای موجود خواهد شد. هال [۱۶] پیشنهاد کرد که $B=1000$ انتخاب گردد تا از همگرایی بازه‌های اطمینان، یقین حاصل شود. شکل ۹ یک نمایش گرافیکی مربوط به یک حلقه از فرآیند بوت استرپ است.

★ : داده های مربوط به نمونه ی ساختگی

سیمار- ویلسون [۴] برای اولین بار از بوت استرپ برای تحلیل حساسیت و ساخت بازه‌های اطمینان برای مقادیر کارایی استفاده کردند. برای این کار آن‌ها داده‌های استفاده شده توسط فار و همکارانش [۱۷] را که این داده‌ها از ۱۹ عامل تأسیسات برقی در سال ۱۹۸۷ گرفته شده بود؛ مورد استفاده قرار دادند. داده‌ها شامل اطلاعاتی در خصوص یک خروجی (قدرت برق بر حسب KWH) و سه ورودی (کارگران بر حسب میانگین استخدام سالانه، سوخت و سرمایه) هستند.

جدول ۱. نتایج مربوط به فرآیند بوت استرپ [۴]

						۲/۵%	۹۷/۵%	۲/۵%	۹۷/۵%
K	$\hat{\theta}_k$	$\bar{\theta}_k$	$Bias_k$	Median of $\hat{\theta}_k^*$	Std. Dev.	Bias Corrected	Centered on $\hat{\theta}_k$		
۱	۸۶۹۲/۰	۸۵۱۹/۰	۰/۱۷۳/۰	۸۴۸/۰	۰/۱۴۳/۰	۸۸۵۴/۰	۰/۸۳۶۰	۰/۸۳۷۲	۰/۹۰۵۷
۲	۱/۰۰۰۰	۰/۹۳۰۷	۰/۰۶۹۳	۰/۹۱۴۵	۰/۰۶۱۴	۱/۰۵۶۴	۰/۸۶۳۱	۰/۸۶۴۶	۱/۰۸۰۰
۳	۱/۰۰۰۰	۰/۹۴۵۷	۰/۰۵۴۳	۰/۹۳۹۶	۰/۰۴۷۵	۱/۰۵۷۴	۰/۸۹۳۲	۰/۸۹۷۳	۱/۰۶۴۷
۴	۰/۹۳۷۰	۰/۹۲۳۷	۰/۰۱۳۳	۰/۹۱۴۸	۰/۰۰۹۴	۰/۹۴۰۰	۰/۹۰۵۹	۰/۹۰۷۶	۰/۹۵۳۶
۵	۱/۰۰۰۰	۰/۹۳۴۹	۰/۰۶۵۱	۰/۹۱۷۷	۰/۰۵۷۳	۱/۰۵۷۳	۰/۸۷۱۷	۰/۸۷۳۰	۱/۰۸۰۷
۶	۰/۹۰۷۱	۰/۸۹۰۶	۰/۰۱۶۵	۰/۸۸۴۹	۰/۰۱۵۱	۰/۹۲۸۶	۰/۸۷۵۶	۰/۸۷۷۷	۰/۹۵۵۲
۷	۰/۸۹۱۵	۰/۸۷۵۹	۰/۰۱۵۶	۰/۸۶۹۳	۰/۰۱۸۵	۰/۹۳۵۳	۰/۸۶۱۵	۰/۸۶۵۲	۰/۹۸۶۱
۸	۰/۸۲۱۰	۰/۸۰۷۵	۰/۰۱۳۵	۰/۸۰۴۱	۰/۰۱۱۸	۰/۸۳۸۷	۰/۷۹۵۵	۰/۷۹۶۷	۰/۸۶۳۰
۹	۰/۸۸۹۲	۰/۸۶۲۴	۰/۰۲۶۸	۰/۸۴۸۸	۰/۰۳۰۱	۰/۹۴۳۴	۰/۸۳۷۰	۰/۸۴۰۵	۰/۹۸۱۶
۱۰	۰/۸۴۶۹	۰/۸۳۷۴	۰/۰۰۹۵	۰/۸۳۵۹	۰/۰۰۶۴	۰/۸۵۴۱	۰/۸۲۹۴	۰/۸۳۰۴	۰/۸۵۹۸
۱۱	۰/۹۵۳۴	۰/۹۴۲۳	۰/۰۱۱۱	۰/۹۳۹۶	۰/۰۱۰۴	۰/۹۷۲۰	۰/۹۳۲۵	۰/۹۳۳۹	۳۷۶۳۱
۱۲	۱/۰۰۰۰	۰/۹۳۳۵	۰/۰۶۶۵	۰/۹۱۵۳	۰/۰۵۹۱	۱/۰۴۸۴	۰/۸۶۸۹	۰/۸۷۰۰	۱/۰۷۴۹
۱۳	۰/۹۶۰۲	۰/۹۴۳۴	۰/۰۱۶۸	۰/۹۳۸۹	۰/۰۱۴۰	۰/۹۷۶۴	۰/۹۲۸۲	۰/۹۳۰۴	۱/۰۰۴۹
۱۴	۱/۰۰۰۰	۰/۹۲۵۸	۰/۰۷۴۲	۰/۹۰۳۶	۰/۰۶۷۶	۱/۰۷۸۶	۰/۸۵۳۴	۰/۸۵۴۵	۱/۰۸۷۲
۱۵	۱/۰۰۰۰	۰/۹۳۳۴	۰/۰۶۶۶	۰/۹۰۸۵	۰/۰۶۴۰	۱/۰۷۸۶	۰/۸۶۸۳	۰/۸۶۹۷	۱/۱۰۲۹
۱۶	۰/۸۸۸۵	۰/۸۷۶۸	۰/۰۱۱۷	۰/۸۷۴۲	۰/۰۰۹۱	۰/۹۰۲۲	۰/۸۶۶۴	۰/۸۶۸۴	۰/۹۱۶۹
۱۷	۱/۰۰۰۰	۰/۹۳۷۸	۰/۰۶۲۲	۰/۹۱۷۵	۰/۰۵۶۱	۱/۰۶۰۸	۰/۸۷۷۴	۰/۸۷۸۳	۱/۰۷۷۴
۱۸	۱/۰۰۰۰	۰/۹۴۲۴	۰/۰۵۷۶	۰/۹۳۲۵	۰/۰۴۹۹	۱/۰۵۲۷	۰/۸۸۶۶	۰/۸۸۷۱	۱/۰۶۴۶
۱۹	۰/۹۴۴۱	۰/۹۳۲۸	۰/۰۱۱۳	۰/۹۳۰۵	۰/۰۰۸۳	۰/۹۵۴۶	۰/۹۲۲۸	۰/۹۲۳۸	۰/۹۶۴۵

نتایجی که ما با اجرای فرآیند بوت استرپ روی داده های فوق به دست آوردیم؛ با صرف نظر از اختلاف ناچیز آنها، می توان گفت که دقیقاً همان نتایج مربوط به سیمار- ویلسون [۴] می باشد. این اختلاف در نتایج، حاصل استفاده از فرآیند انتخاب تصادفی از بین داده ها و توزیع نرمال استاندارد است. جدول ۱ نتایج مربوط به آزمایش بوت استرپ را با $B=1000$ و $h=0/014$ نشان می دهد. ستون ۱ نمایش مربوط به واحدها است که تعداد آنها ۱۹ واحد می باشد و ستون های ۲ تا ۶ مقادیر برآورد کارایی DEA اصلی، برآورد با اریب اصلاح شده، برآورد اریب بوت استرپ، مقادیر میانگین بوت استرپ و انحراف استاندارد آنها را نمایش می دهد. چهار ستون آخر یک بازه اطمینان ۹۵٪ را نمایش می دهد. بازه اطمینان اولی بر اساس اصلاح مقدار اریبی که در (۲۳) آمده و بازه دومی با اصلاح مقدار اریبی از بازه های اولیه به دست آمده است. از آنجا که مقادیر اریبی، مقادیر کوچکی هستند؛ دو بازه به دست آمده بسیار شبیه به هم است.

نتایج ثبت شده در جدول ۱ حساسیت اندازه کارایی را نسبت به تعدد نمونه گیری آشکار می کند. با توجه به نتایج جدول ۱ این نکته را در نظر بگیریم که در رتبه بندی نسبی واحدها بر حسب مقادیر کارایی اصلی باید کمی با دقت عمل کنیم. برای مثال واحد یک دارای کارایی DEA، $\hat{\theta}_1 = 0/8962$ ، در حالی که واحد ۲ دارای کارایی DEA $\hat{\theta}_2 = 1.0$ است. با توجه به مقادیر اندازه کارایی اریب- اصلاح شده ستون ۳ مشاهده می کنیم که اختلاف بسیار ناچیز است، اما به هر حال بین آنها تفاوتی وجود دارد. در چهار ستون آخر که نمایش دهنده ی بازه ی اطمینان برای مقادیر کارایی است نقاط اشتراک زیادی بین بازه های اطمینان واحد ۱ و بازه های اطمینان واحد ۲ وجود دارد. بنابراین نمی توان گفت که این دو واحد به خاطر داشتن مقادیر کارایی DEA اصلی متمایز، آشکارا متمایزند (از لحاظ کارایی).

۳-۴ چارچوب کلی برای شبیه سازی در زمینه ی مرزهای ناپارامتری

متدولوژی ارایه شده در بخش قبل، با توجه به فرض های محدود کننده ای به دست آمده بود (فرض همگن بودن توزیع مقادیر کارایی). سیمار- ویلسون [۱۳] یک چارچوب کلی برای پیاده سازی بوت استرپ، در چارچوب مدل های مرزی ناپارامتری ارایه کردند. آنها فرض های متعددی روی DGP بنا نمودند که حالت عمومی تری نسبت به قبل داشت.

مدل آماری معرفی شده توسط سیمار و ویلسون، شامل فرض های زیر است که توسط نیپ [۱۸] و همکارانش استفاده شده بود. الگوریتم به صورت زیر است:

$$۱- \text{مقادیر کارایی } \hat{\theta}_k \text{ را برای هر مشاهده و } \hat{\delta}_k = \frac{1}{\hat{\theta}_k} \text{ را محاسبه می کنیم.}$$

۲- با استفاده از روش بازتابی سیلورمن این مقادیر را نسبت به $\hat{\delta}=1$ منعکس کرده تا مجموعه زیر به دست آید.

$$\{\hat{\delta}_1, \hat{\delta}_2, \dots, \hat{\delta}_n, 2-\hat{\delta}_1, 2-\hat{\delta}_2, \dots, 2-\hat{\delta}_n\}$$

۳- یک نمونه تصادفی و با جای گذاری به حجم n از مجموعه به دست آمده در مرحله قبل استخراج می کنیم.

$$\{\hat{\delta}_1^*, \hat{\delta}_2^*, \dots, \hat{\delta}_n^*\}$$

۴- با استفاده از مجموعه مرحله ۳، مجموعه نقاط هموار $\{\tilde{\delta}_1^*, \tilde{\delta}_2^*, \dots, \tilde{\delta}_n^*\}$ را می سازیم.

$$\tilde{\delta}_i^* = \hat{\delta}_i^* + \varepsilon_i^* h$$

که ε_i^* به صورت تصادفی از توزیع نرمال استاندارد انتخاب می شود.

۵- در این مرحله با تصحیح میانگین و واریانس مقادیر هموار مرحله ۴، مقادیر به دست آمده را اصلاح می کنیم.

$$\tilde{\tilde{\delta}}_i^* = \bar{\tilde{\delta}} + \frac{\tilde{\delta}_i^* - \bar{\tilde{\delta}}}{\sqrt{1 + \frac{h^2}{s^{*2}}}}$$

۶- بر خلاف بخش قبل که نیمه سمت چپ تابع برآوردگر $\hat{f}_{s,h}(t)$ ، برآوردگر تابع ناشناخته اصلی بود؛ به دلیل تفاوت در شرط مرزی، در اینجا نیمه سمت راست تابع $\hat{f}_{s,h}(t)$ ، برآوردگر اصلی است. لذا می بایست به نقاط بزرگ تر از ۱ برگردیم. از این رو:

$$\tilde{\tilde{\delta}}_i^* = \begin{cases} \tilde{\tilde{\delta}}_i^* & \tilde{\tilde{\delta}}_i^* < 1 \\ \tilde{\tilde{\delta}}_i^* & \text{O.W} \end{cases}$$

۷- نمونه بوت استرپ به صورت زیر ساخته می شود:

$$\chi_n^* = \left\{ (X_i^*, Y_i) \mid X_i^* = \frac{\tilde{\tilde{\delta}}_i^* X_i}{\tilde{\tilde{\delta}}_i^*}, i = 1, \dots, n \right\}.$$

۸- با استفاده از برآوردگر DEA مقادیر برآوردهای بوت استرپ را برای نمونه ساختگی χ_n^* و واحد دلخواه و ثابت k محاسبه می کنیم؛ یعنی:

$$\hat{\delta}_b^k = \sup \left\{ \theta \mid \frac{X}{\theta} \in \Psi^* \right\}.$$

۹- با تکرار مراحل بالا به تعداد B بار مجموعه $\{\hat{\delta}_b^k\}_{b=1}^B$ به دست می آید.

۱۰- توزیع تجربی $\{\hat{\delta}_b^k\}_{b=1}^B$ تخمین بوت استرپ از توزیع نمونه ای $\hat{\delta}_k$ می باشد.

تنها مطلبی که در این بخش باقی می ماند؛ نحوه انتخاب پارامتر هموارسازی است. برای این کار باید دقت زیادی به خرج داد؛ چرا که اگر مقدار h کوچک انتخاب شود به طوری که در شرایط حدی نزدیک صفر باشد؛ آن گاه تابع $\hat{f}_h(t)$ به تابع چگالی تجربی گسسته و اگر h مقدار بزرگی انتخاب شود؛ به طوری که در شرایط حدی به بی نهایت میل کند؛ آن گاه تابع $\hat{f}_h(t)$ به یک تابع چگالی یکنواخت صاف همگرا می شود. سیلورمن [۵] نشان داد

که اگر $f(\circ)$ گوسی باشد؛ آن گاه مقدار بهینه h ، خطای مربع مجموع میانگین بین $f(\delta)$ و $\hat{f}(\delta)$ را کمینه می کند که به صورت زیر محاسبه می شود:

$$h = 1/0.6 s_n n^{\frac{-1}{5}}. \quad (57)$$

که s_n انحراف استاندارد تجربی n مقدار $\hat{\delta}_k$ هاست و قانون مرجع نرمال نامیده می شود. هم چنین سیلورمن نشان داد که

$$h = 1/0.6 \min(s_n, \frac{r_n}{1/34}) n^{\frac{-1}{5}}, \quad (58)$$

که r_n محدوده ی درون چارک n داده $\hat{\delta}_k$ هاست و قانون مرجع نرمال تنومند نامیده می شود.

۵ نتیجه گیری

تحلیل پوششی داده ها یک روش ناپارامتری برای محاسبه ی اندازه کارایی یک گروه از واحدهایی است که فعالیت یکسانی را انجام می دهند. کارایی به دست آمده از این روش یک اندازه نسبی است و به تغییرات نمونه حساسیت شدیدی نشان می دهد. بوت استرپ که برای اولین بار توسط سیمار در مدل های تحلیل پوششی داده ها مورد استفاده قرار گرفت؛ ابزار سودمندی برای بهبود مقادیر کارایی به دست آمده از این روش می باشد. این فرآیند ضمن ساخت بازه های اطمینان برای اندازه کارایی و مقادیر کارایی بهبود یافته قدری از شرایط عدم اطمینان کاسته و شرایط بهتری برای مقایسه و رتبه بندی واحدها فراهم می کند. هم چنین بازه های اطمینان میزان حساسیت کارایی هر واحد را به تغییر نمونه نشان می دهد. هرچه طول بازه کارایی کمتر باشد؛ حساسیت اندازه کارایی به مدل کمتر بوده؛ میزان اعتماد به مقدار کارایی به دست آمده برای واحد بیشتر خواهد بود. از بازه های اطمینان می توان برای رتبه بندی واحدها نیز استفاده نمود. برای این کار باید به برخی موارد دقت کرد. طول بازه اطمینان (در مواردی که کارایی به دست آمده از DEA برای دو واحد متفاوت نزدیک به هم باشد)، محدوده بازه اطمینان و میزان اشتراک بازه ها اطمینان، واحدهایی که با یکدیگر مقایسه می شوند.

تحلیل پوششی داده ها ابزار قدرتمندی برای محاسبه کارایی است. با این وجود حساسیت مدل های آن به تغییر نمونه سبب کاهش اطمینان به مقادیر کارایی های به دست آمده از این روش می شود (در مواردی که حجم نمونه زیاد باشد؛ مرز کارایی به دست آمده، تنومندتر بوده؛ نسبت به تغییر نمونه حساسیت کمتری دارد. لذا استفاده از بوت استرپ برای چنین مواردی توصیه نمی شود). در چنین مواردی پیشنهاد می شود فرآیند شبیه سازی بوت استرپ، برای تحلیل میزان حساسیت کارایی به تغییر نمونه، مورد استفاده قرار گیرد.

منابع

- [۲] عبادی، س.، (۱۳۹۰). روشی برای رتبه‌بندی نمرات کارایی با استفاده از بوت استرپ، مجله ریاضیات کاربردی واحد لاهیجان، ۲۸(۲)، ۴۴-۲۹.
- [1] Efron, B. (1979), Bootstrap methods: another look at the jackknife, *Annals of Statistics* 7, 1-16.
- [3] Ferrier, G. D. and J.G. Hirschberg (1997), Bootstrapping confidence intervals for linear programming efficiency scores: With an illustration using Italian bank data, *Journal of Productivity Analysis* 8, 19-33.
- [4] Simar, L. and Wilson, P.W. (1998a), "Sensitivity analysis of efficiency scores: how to bootstrap in nonparametric frontier models", *Management Science*, vol. 44, 1, 49-61.
- [5] Banker, R.D. (1993), Maximum likelihood, consistency and data envelopment analysis: A statistical foundation, *Management Science* 39, 1265-1273.
- [6] Banker, R.D. (1996), Hypothesis tests using data envelopment analysis, *Journal of Productivity Analysis* 7, 139-159.
- [7] Korostelev, A., L. Simar, and A.B. Tsybakov (1995a), efficient estimation of monotone boundaries, *The Annals of Statistics* 23, 476-489.
- [8] Korostelev, A., L. Simar, and A.B. Tsybakov (1995b), On estimation of monotone and convex boundaries, *Pub. Inst. Stat. Univ. Paris*, XXXIX, 18-31.
- [9] Advanced robust and non-parametric method in efficiency analysis; methodology and applications, Cinzia Daraio and Leopold Simar (2007).
- [10] Efron, B. and R.J. Tibshirani (1993), *An Introduction to the Bootstrap*, London: Chapman and Hall.
- [11] S.C., Ray (2004), *Data Envelopment Analysis Theory and Techniques for Economics and Operations Research*, university of Connecticut.
- [12] Simar, L. (1992), Estimating efficiencies from frontier models with panel data: A comparison of parametric, non-parametric and semi-parametric methods with bootstrapping, *Journal of Productivity Analysis* 3, 167-203.
- [13] Simar, L. and Wilson, P.W. (2000b), "A general methodology for bootstrapping in non-parametric frontier models", *Journal of Applied Statistics*, vol. 27, 6, 779-802.
- [14] Silverman, B.W. (1986), *Density Estimation for Statistics and Data Analysis*, London: Chapman and Hall.
- [15] Schuster, E. F.(1985), Incorporating support constraint into nonparametric estimate of densities, *communications in statistics- theory and method*, 14, 1123-1136.
- [16] Hall, P.(1986), On the number of bootstrap simulations required to construct a confidence interval , *Ann. Statistics*, 14, 1453-1462.
- [17] Fare, R., Grosskopf, S., (1985). A nonparametric cost approach to scale efficiency. *Scandinavian Journal of Economics* 87, 594-604.
- [18] Kneip, A., L. Simar, and P.W. Wilson (2003), "Asymptotics for DEA Estimators in Nonparametric Frontier Models," Discussion paper #0317, Institut de Statistique, Université Catholique de Louvain, Louvain-la-Neuve, Belgium.
- [19] Simar, L. and Wilson, P.W. (1999b), "Of Course we Can Bootstrap DEA scores! But does it mean anything? Logic Trumps and Wishful Thinking", *The Journal of Productivity Analysis*, 11, 67-80.