

خوشه‌بندی خودروسازان بر اساس تابع تولید اقتصادی با استفاده از تحلیل پوششی داده‌ها

سمیه رضایی^۱، دکتر غلام رضا امین^۲، دکتر میربهادرقلی آریان‌نژاد^{۳*}

چکیده

خوشه‌بندی به کمک تحلیل پوششی داده‌ها (DEA) شناسایی روابط پنهان بین عوامل ورودی و خروجی واحدهای تصمیم‌گیری در تعیین تابع تولید آنهاست. در این مقاله به کمک خوشه‌بندی بر اساس DEA توابع تولید واحدهای تصمیم‌گیری صنایع خودروسازی (از جمله سایپا) به صورت تفکیک شده مشخص می‌شود. تعیین توابع تولید خودروسازان و استفاده از آنها در تفکیک صنایع مشابه با قابلیت در نظر گرفتن همزمان چندین عامل ورودی و خروجی از مزیت‌های روش خوشه‌بندی بر اساس DEA است. در نتیجه به‌کارگیری این روش نه تنها دسته‌ای را که هر واحد خودروسازی به آن تعلق دارد مشخص شده است، بلکه نوع تابع تولید واحدهای هر دسته نیز شناسایی می‌گردد. از طرف دیگر با مقایسه تابع تولید به‌کاررفته در خوشه‌های متفاوت اطلاعات مهمی در مورد چگونگی استفاده از منابع در ترکیب ورودی‌ها برای صنایع متفاوت به‌دست می‌آید.

واژه‌های کلیدی: تحلیل پوششی داده‌ها، خوشه‌بندی، تابع تولید، صنایع خودروسازی، سایپا

Automakers Clustering based on Economic Production Function using Data Envelopment Analysis

S. Rezaei, M.Sc; Gh.R. Amin, Ph. D.; M.Gh. Ariyannezhad, Ph. D.

Abstract

DEA-based clustering approach reveals the input-output relationships hidden in the data items of input and output. DEA-based clustering approach employs the piece-wise production functions derived from the DEA method to cluster the car companies. Estimate production function for each car company by input-output data is the benefit of this method. Thus, each car company (like Saipa) not only knows the cluster that it belongs to, but also checks the production function type that it confronts. It is important for managerial decision-making in different fields where decision-makers are interested in knowing the changes required in combining input resources.

Keywords: Data Envelopment Analysis, Clustering, Production Function, Automakers, Saipa

۱. کارشناس ارشد مدیریت صنعتی گرایش تحقیق در عملیات، دانشگاه آزاد اسلامی - واحد تهران جنوب

۲. دانشیار، عضو هیئت علمی دانشگاه آزاد اسلامی واحد تهران جنوب

۳. استاد، عضو هیئت علمی دانشگاه علم و صنعت ایران

*پست الکترونیکی نویسنده اصلی: rezaei.somaye@gmail.com

مقدمه

خوشه‌بندی در واقع تقسیم‌بندی یک جمعیت ناهمگون^۱ به تعدادی از زیرمجموعه‌های همگون^۲ است که به آنها خوشه اطلاق می‌شود. هدف از خوشه‌بندی یافتن گروه‌هایی است که با یکدیگر بسیار متفاوت‌اند ولی اعضای درون این گروه‌ها بسیار شبیه به هم هستند. به عبارت دیگر در خوشه‌بندی، مجموعه‌هایی از عناصر یا افراد را بر مبنای مشابهت آنها به چند زیر مجموعه کوچک‌تر تقسیم می‌کنیم، به طوری که بین عناصر در یک خوشه بیشترین شباهت و بین عناصر واقع در دو خوشه مجزا بیشترین تمایز وجود داشته باشد. از متداول‌ترین روش‌های خوشه‌بندی می‌توان به روش‌های خوشه‌بندی سلسله‌مراتبی، خوشه‌بندی شبکه‌ای، خوشه‌بندی بر اساس تابع هدف و خوشه‌بندی بر اساس بخش‌بندی اشاره کرد. اغلب این الگوریتم‌ها روش‌های مبتنی بر فاصله هستند، الگوریتم K-means و الگوریتم فازی C-means (FCM) نمونه‌هایی از این الگوریتم‌ها هستند. از زمانی که این مسئله مورد توجه محققین قرار گرفته است ما شاهد تحقیقات بسیاری در زمینه تکنیک‌های خوشه‌بندی هستیم. با این وجود گاهی نتایج به دست آمده از این روش‌ها پاسخ‌گوی نیازهای مدیران در صنایع مختلف نیستند. به عنوان مثال یکی از مهمترین نیازهای مدیران به خصوص در واحدهای تولیدی شناسایی شباهت‌های موجود میان توابع تولید واحدهای مختلف به منظور چگونگی استفاده از منابع در ترکیب ورودی‌ها است. خوشه‌بندی بر اساس تحلیل پوششی داده‌ها با استفاده از توابع به دست آمده از داده‌های ورودی و خروجی می‌تواند برطرف کننده این نیاز برای صنایع مختلف و همچنین صنعت خودروسازی باشد. در نتیجه به کارگیری این روش نه تنها دسته‌ای را که هر واحد خودروسازی بدان تعلق دارد مشخص شده، بلکه نوع تابع تولید واحدهای هر دسته نیز شناسایی می‌گردد.

در این پژوهش ضمن بیان مختصری از مدل CCR و همچنین خصوصیات تابع تولید، الگوریتم خوشه‌بندی بر اساس DEA ارائه و سپس به پیاده‌سازی این مدل در شرکت‌های خودروسازی و بررسی نتایج این مدل می‌پردازیم.

تحلیل پوششی داده‌ها

در دهه‌های اخیر شاهد توسعه و گسترش قابل ملاحظه‌ای در

پیشبرد تئوری‌های ریاضی بهینه‌سازی برای حل مسائل پیچیده تصمیم‌گیری بوده‌ایم. در این بین تحلیل پوششی داده‌ها از جذاب‌ترین و معروف‌ترین مدل‌های برنامه‌ریزی خطی هستند که در عمل به کار گرفته شده‌اند و برای مقایسه، ارزیابی، تعیین کارایی و انتخاب واحدهای تصمیم‌گیری راهکارهای ارزشمندی را ارائه داده‌اند. جذابیت عنوان شده در DEA توسط مدیرانی که آن را برای مسائل تصمیم‌گیری به کار گرفته‌اند نشان داده شده است. تحلیل پوششی داده‌ها، یک تکنیک برنامه‌ریزی ریاضی^۳ است که کارایی نسبی گروهی از واحدهای تصمیم‌گیری^۴ (DMUs) را اندازه‌گیری می‌کند یا به عبارت دیگر، DEA یک برنامه‌ریزی ریاضی جهت اندازه‌گیری واحدهای سازمانی که دارای نهاده‌ها و ستاده‌های مختلف بوده و کار مقایسه و سنجش کارایی مشکل است، می‌باشد. تحلیل پوششی داده‌ها یک روش غیر پارامتریک بوده که به کمک برنامه‌ریزی ریاضی به تعیین مرز کارایی واحدهای تصمیم‌گیری که دارای ستاده‌ها و نهاده‌های مشابه‌اند می‌پردازد.

در حالت کلی فرض کنیم که تعداد n مشاهده یا واحد تصمیم‌گیری در دسترس داریم. همچنین فرض کنید که مشاهده j ام با مصرف بردار m تائی ورودی $x_j = (x_{1j}, \dots, x_{mj})^t$ بردار s تائی خروجی $y_j = (y_{1j}, \dots, y_{sj})^t$ را تولید نموده است که در آن $j = 1, \dots, n$ می‌باشد. برای واحد k ام ورودی و خروجی مجازی^۵ زیر را تعریف می‌نمائیم و کارایی این واحد عبارت خواهد بود از:

$$\max V_k / U_k = \max \frac{\sum_{i=1}^m v_i x_{ik}}{\sum_{r=1}^s u_r y_{rk}}$$

مدل CCR را می‌توان به یک بازار تشبیه کرد که در آن واحد تحت ارزیابی k می‌تواند ورودی‌هایش را به هر قیمتی بخرد و خروجی‌هایش را به هر قیمتی بفروشد، اما واحدهای دیگر نیز می‌توانند ورودی و خروجی‌هایشان را به قیمت‌هایی که واحد k معامله می‌کند معامله کنند. حال در این بازار رقابتی واحدی کارا است که نسبت میزان فروش به میزان خرید بیشتر شود. لازم به ذکر است که مدل CCR بیان شده در فوق کسری است و در عمل از خطی‌شدن آن استفاده می‌شود:

1. Heterogeneous population
2. Homogeneous
3. Linear programming

4. Decision Making Units
5. Virtual input and output

اصل ۲. ناکارا بودن^۳:

الف: چنانچه $(x, y) \in T$ و $\bar{x} \succ x$ باشد
آنگاه $(\bar{x}, y) \in T$ است.

ب: چنانچه $(x, y) \in T$ و $\bar{y} \prec y$ باشد
آنگاه $(x, \bar{y}) \in T$ است.

اصل ۳. نامحدود بودن شعاع^۴: اگر $(x, y) \in T$ باشد آنگاه
برای هر $\lambda > 0$ داریم $(\lambda x, \lambda y) \in T$. با توجه به شکل (۱)
کلیه نقاط روی مرز کارا و ضرایب آنها امکان تولید دارند.

اصل ۴. غیر تهی بودن^۵: به ازای
هر $(x_j, y_j) \in T, (j = 1, 2, \dots, m)$ ، ویژگی غیر تهی
بودن را ویژگی شمول مشاهدات نیز می‌گویند و به معنی آن است
که تمام مشاهدات در T (در شکل (۱)) با PPS نشان داده شده
است (قرار دارند).

اصل ۵. کمینه برون‌یابی^۶: T کوچکترین مجموعه‌ای است
که ویژگی‌های ۱، ۲، ۳، ۴ را دارا است. مجموعه امکان تولید
منحصر به فرد T_c با اصول ۱، ۲، ۳، ۴ با رابطه زیر مشخص
می‌شود.

$$T_c = \{(x, y) | x \geq \sum_{j=1}^n \lambda_j x_j, y \geq \sum_{j=1}^n \lambda_j y_j, \lambda_j \geq 0, j=1, 2, \dots, n\}$$

که مرز این مجموعه خط مستقیمی است که از مبدا می‌گذرد،
به گونه‌ای که T_c یک مخروط محدب است که در برگرفته تمام
واحدها می‌باشد که به آن مجموعه امکان تولید CCR می‌گویند.
(شکل ۱).

خوشه‌بندی به کمک تحلیل پوششی داده‌ها

خوشه‌بندی به کمک تحلیل پوششی داده‌ها روشی کاملاً
جدید است. روش خوشه‌بندی بر اساس تحلیل پوششی داده‌ها را
با استفاده از مثال زیر توضیح می‌دهیم. همان‌طور که در شکل
مشاهده می‌شود روش DEA از یک تابع تولید قطعه به قطعه
که به عنوان مرز کارایی فارل معرفی شده است استفاده می‌کند
که این مرز توسط واحدهای تصمیم‌گیری (DMU_1, DMU_2)
 (DMU_3, DMU_4) و پنج تابع تولید با ضرایب مجازی مختلف
۱ ایجاد می‌شود. این توابع تولید با توجه به رابطه زیر به دست
می‌آیند:

$$\begin{aligned} \text{Max} \quad & \sum_{r=1}^s u_r y_{rp} \\ \text{st} \quad & \sum_{i=1}^m v_i x_{ip} = 1, \\ & \sum_{r=1}^s u_r y_{rj} - \sum_{i=1}^m v_i x_{ij} \leq 0, \quad j=1, \dots, n \\ & u_r \geq \varepsilon, \quad r=1, \dots, s \\ & v_i \geq \varepsilon, \quad i=1, \dots, m \end{aligned}$$

مدل CCR با ماهیت ورودی است. یعنی مرز کارا را در جهت
کمینه‌سازی ورودی تا جایی که خروجی کاهش نیابد محاسبه می-
کند. نظیر همین مدل را با ماهیت خروجی می‌توان نوشت.
مدل‌های با ماهیت خروجی مرز کارا را در جهت حداکثر کردن
خروجی تا جایی که ورودی افزایش نیابد محاسبه می‌کنند.

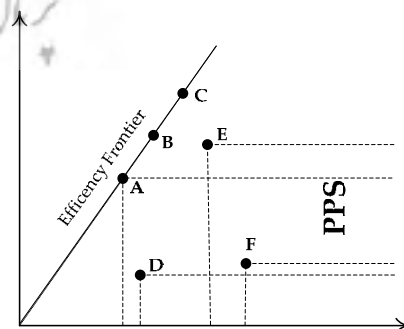
مجموعه امکان تولید^۱ (PPS)

مجموعه امکان تولید n ، $j = (1, 2, \dots, n)$ واحد با m
ورودی مختلف $x_j = \{x_{1j}, x_{2j}, \dots, x_{mj}\}$ و s خروجی
 $y_j = \{y_{1j}, y_{2j}, \dots, y_{sj}\}$ به صورت زیر و با نماد T تعریف می‌شود.
 $T = \{(X, Y) | X \text{ ایجاد شود توسط } Y\}$
مجموعه امکان تولید T به طور تجربی به وسیله اصول زیر
تعریف می‌شود.

اصل ۱. محدب بودن^۲: چنانچه

$(x_j, y_j) \in T, (j = 1, 2, \dots, m)$ و $\lambda_j \geq 0$ به گونه‌ای

که $\sum_{j=1}^n \lambda_j = 1$ باشد، آنگاه $(\sum_{j=1}^n \lambda_j x_j, \sum_{j=1}^n \lambda_j y_j) \in T$
می‌باشد. در شکل (۱) کلیه نقاط زیر مرز کارا امکان تولید دارند.



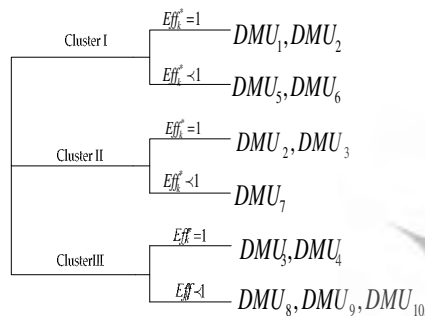
شکل ۱: مجموعه امکان تولید CCR

1. Production Possibility Set
2. Convexity
3. Inefficiency

4. Ray unboundness
5. Inclosing of observation
6. Minimum Extrapolation

Cluster V : $(PF(5): 1 = v_{52}^* x_2) \{DMU_1\}$

همان‌طور که مشاهده می‌شود، تمام ضرایب (v_{ij}^*) توابع تولید مطابق با هر یک از دسته‌های I، II و III غیر صفر هستند (دسته‌های موثر^۲)، اما در بین خوشه‌ها دسته‌های IV و V وجود دارند که بعضی از ضرایب تابع تولید در آنها صفر است (دسته‌های تباهیده^۳)، این توابع نمی‌توانند به‌عنوان یک خوشه در نظر گرفته شوند. بنابراین واحدهای ناکارایی (مانند DMU₁₀) که در این خوشه‌ها وجود دارند باید دوباره خوشه‌بندی شوند. بر اساس اصل کمینه برون‌یابی از مجموعه امکان تولید نزدیکترین تابع تولید به هر یک از این واحدهای ناکارا باید به‌عنوان خوشه آن واحد در نظر گرفته شود، در نتیجه پنج خوشه تولید شده در مرحله اول به سه خوشه زیر کاهش می‌یابد.



برای دسته‌بندی واحدهای تباهیده (مانند DMU₁₀) به نزدیک‌ترین دسته ابتدا ورودی‌های واحد مورد نظر را در $t(i)$ ضرب می‌کنیم (i نشانگر دسته مورد نظر است). در مرحله بعد با قرار دادن هر یک از توابع تولید $(PF(i))$ مساوی صفر و جایگزین کردن مقادیر (x, y) مقدار $t(i)$ را به ازای هر یک از توابع به‌دست می‌آوریم در مرحله آخر، ماکزیم مقدار $t(i)$ تعیین‌کننده نزدیکترین خوشه‌ای است که واحد تباهیده باید به آن دسته‌بندی شود.

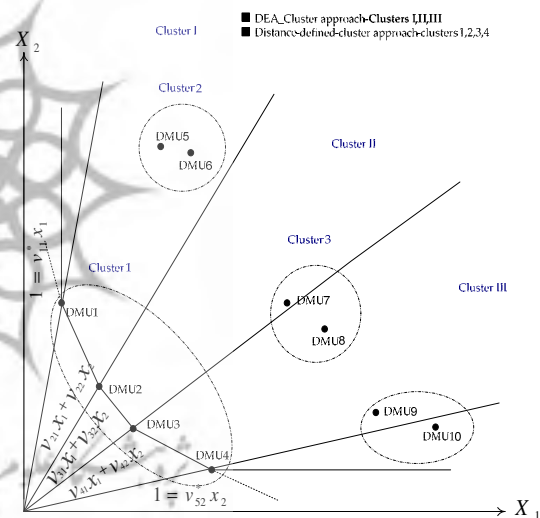
الگوریتم خوشه‌بندی بر اساس DEA

از آنجایی که در دنیای واقعی و مسائلی با ابعاد بیشتر قادر به رسم نمودار و شناسایی توابع تولید با استفاده از شکل نیستیم، روش ارائه شده در بالا برای شناسایی توابع تولید و همچنین شناسایی روابط موجود میان توابع تولید ایجاد شده به صورت الگوریتم زیر ارائه شده است. با استفاده از این الگوریتم همزمان دسته‌های موثر و تباهیده شناسایی و از هم تفکیک می‌شوند.

$$f(x_1, x_2, \dots, x_m, y_1, y_2, \dots, y_s) = \sum_{r=1}^s u_r^* y_{rk} - \sum_{i=1}^m v_i^* y_{ik} = 0$$

بنابراین توابع تولید مربوط به قطعات ایجاد شده بر روی مرز کارا عبارت خواهند بود از :

$1 = v_{11}^* x_1, 1 = v_{21}^* x_1 + v_{22}^* x_2, 1 = v_{31}^* x_1 + v_{32}^* x_2, 1 = v_{41}^* x_1 + v_{42}^* x_2, 1 = v_{52}^* x_2$ توجه شود که در تابع پنج تولید به‌دست‌آمده در بالا مقدار خروجی یا y برای کلیه واحدها برابر یک است. طبق نظریه فارل، هر واحد تصمیم‌گیری ناکارا را می‌توان برحسب تعدادی از واحدهای کارا نوشت که این واحدهای کارا مجموعه مرجع^۱ نام دارند. (به عنوان مثال با توجه به شکل ۲ واحد تصمیم‌گیری ۶ ناکاراست و برحسب واحدهای کارای ۱ و ۲ قابل نوشتن است بنابراین واحدهای ۱ و ۲ مجموعه مرجع‌های واحد ۶ هستند، پس تمامی واحدهای ناکارا دارای مراجعی بر روی مرز کارا بوده و دارای توابع تولیدی مشابه یکی از این ۵ تابع تولید به‌دست‌آمده می‌باشند. در نتیجه می‌توان تمامی واحدهای تصمیم‌گیری را به پنج دسته خوشه‌بندی کرد.



شکل ۲: نتایج به دست آمده از خوشه‌بندی به روش DEA و روشهای مبنی بر فاصله

Cluster I : $(PF(1): 1 = v_{21}^* x_1 + v_{22}^* x_2)$

$\{DMU_1, DMU_2, DMU_5, DMU_6\}$

Cluster II : $(PF(2): 1 = v_{31}^* x_1 + v_{32}^* x_2)$

$\{DMU_2, DMU_3, DMU_7\}$

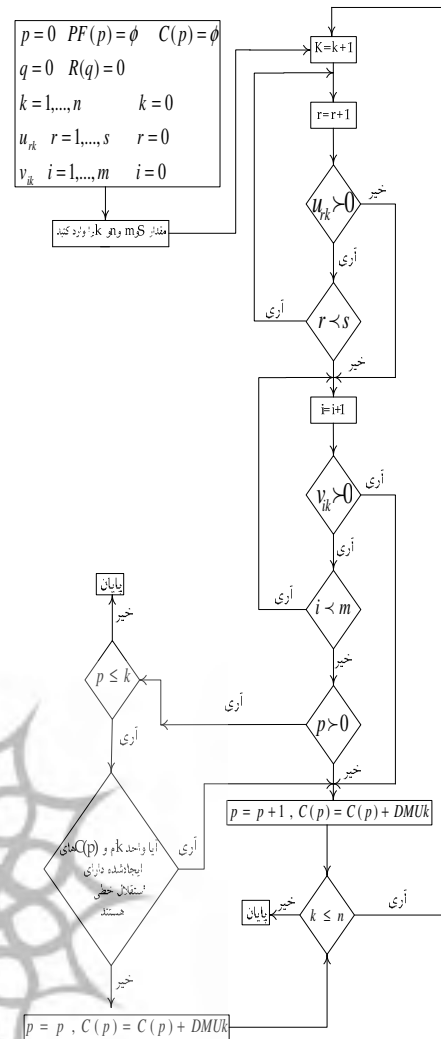
Cluster III: $(PF(3): 1 = v_{41}^* x_1 + v_{42}^* x_2)$

$\{DMU_3, DMU_4, DMU_8, DMU_9\}$

Cluster IV: $(PF(4): 1 = v_{11}^* x_1) \quad \{\}$

جدول ۱: ۷۶ شاخص‌های موجود

Revenue	Tangible Assets	Op. Profit / Unit Sold	Other Current Assets
Cost Of Sales	Other Non-Current	Op. Profit / Revenue	Capital Investment
Gross Profit	Total Non-Current Assets	Op. Profit Per Employee	Automobile revenue
Materials	Total Assets	Op. Profit / Cap. Employed	Raw Materials
Labour	Trade Creditors	PBT / Unit Sold	Work In Progress
Depreciation	Other Current Liabilities	PBT / Revenue	Finished Goods
Other Operating Expense	Total Current Liabilities	PBT / Employee	Total Inventory
Other Operating Income	Provisions / Other	PBT / Capital Employed	Production / Employee
Total Operating Expense	Long Term Debt	Tax / PBT	Material / Unit Produced
Operating Profit / (Loss)	Shareholders' Equity	Net Profit / Unit Sold	Material / Op. Exp.
Interest Income	Total Liabilities	Net Profit / Revenue	Labour / Unit Produced
Interest Expense	Average Employees	Value Added	Trade Debtors
Other Income / (Expense)	Production	Average Inventory Level	Cash
Profit / (Loss) Before Tax	Sales	R&D / Revenue	Interest Cover
(Tax) / Tax Credit	Revenue change	Capex / Revenue	Liquidity (Acid Ratio)
Other Items	Revenue Per Unit Sold	Capex / Cash Flow	Equity / Total Assets
Net Profit / (Loss)	Auto. Revenue / Unit Sold	Average Debtor Turnover	Op. Profit change
Cash Flow	Revenue / Employee	Average Creditor Turnover	Total Current Assets
R&D	Revenue / Cap. Employed	C. Assets / C. Liabilities	Labour / Total Op. Exp.



خوشه‌بندی خودروسازان

اطلاعات و داده‌های مربوط به این پژوهش با مراجعه به اسناد و مدارک خودروسازی سایپا به‌دست‌آمده است. در این مقاله به‌منظور تعیین واحدهای تصمیم‌گیری و همچنین انتخاب متغیرهای ورودی و خروجی از اطلاعات موجود در مورد شرکتهای خودروسازی استفاده شده است. گزارش‌های سالیانه حاوی ۷۶ شاخص برای ۵۹ شرکت خودروسازی و برای ۲۰ سال اخیر مورد بررسی قرار گرفته و سعی در استفاده از جدیدترین داده‌ها شده است. این شاخص‌ها و شرکت‌های خودروسازی در جدول ۱ ارائه شده‌اند.

جدول ۲: ۵۹ شرکت خوروسازی

BMW AG	Ford Europe	Hyundai unconsolidated	Renault Auto
BMW Group	Ford Group	Hyundai consolidated	Renault Group
GM Daewoo Auto & Technology	Ford UK	Kia	Suzuki
Chrysler	Ford Werke	Isuzu non-cons.	Suzuki non-cons
Chrysler (CFC equity basis)	Jaguar	Izuzu cons.	Tata
Chrysler Car & Truck	Volvo Car	Mazda cons.	Daihatsu
D-Benz	Fuji Cons.	Mazda non-cons.	Toyota
DC	Fuji non-Cons.	MG Rover	Volvo Group
DC Auto	GM	Mitsubishi Motors	Audi
M-Benz	GM Auto	MMC Auto	Seat
M-Benz Cars	GM Europe	Nissan	Skoda
FAW	GM GMAC equity basis	Porsche	VW Group
Fiat Auto	Saab	Proton	VW Automotive
Fiat Group	Honda	PSA	Saipa
Ford Auto	Honda Auto	PSA Auto	

متغیرهای دیگری مانند بدهی‌های بلندمدت، متغیرهای مربوط به سهام و منابع مالی بیشتر مطرح می‌شد.

همچنین با توجه به آنکه متغیرهای مورد نظر برای تمام شرکت‌ها، برای سال‌های مختلف، در اختیار نبود لازم بود برخی شرکت‌ها حذف شوند.

متغیرهای پژوهش با توجه به روش استفاده شده شامل موارد زیر می‌باشد:

عوامل ورودی^۲:

تحقیق و توسعه (Research & Development; R&D)

دارایی کل (Total Assets)

متوسط کارکنان مشغول به کار (Average Employment)

عوامل خروجی^۳:

تولید (Production)

ارزش افزوده (valu added)

درآمد (Revenue)

سود خالص (Net Profit)

در اینجا روش ارائه شده در بخش قبل را بر روی ۲۳ شرکت خودروسازی است اجرا کرده و به بررسی نتایج آن می‌پردازیم. واحدهای خودروسازی همگی مربوط به سال ۲۰۰۶ هستند و توسط شرکت خودروسازی سایپا فراهم شده‌اند.

در نهایت با استفاده از نظر خبرگان سایپا، ۲۳ شرکت خودروسازی از میان ۵۹ شرکت موجود به‌عنوان واحدهای تصمیم‌گیری انتخاب شدند. به‌عنوان مثال شرکت هیوندای^۱ با توجه به اینکه علاوه بر خودرو تولیدکننده محصولات دیگر نیز هست از لیست خودروسازان حذف شد.

متغیرهای مورد بررسی در این مقاله شامل ۳ ورودی و ۴ خروجی است که با توجه به محدودیت‌هایی که در زمینه جمع‌آوری داده‌ها وجود داشت و با توجه به اهمیت برخی از داده‌های تولیدی از دیدگاه خبرگان سایپا و همچنین با توجه به رابطه آماری که بین شاخص‌ها وجود داشت از میان ۷۶ شاخص پیشنهادی انتخاب شدند.

در انتخاب این متغیرها معیارهای زیر در نظر گرفته شده‌اند:

۱- متغیرها از نظر مفهومی از هم مستقل باشند. مثلاً فروش در کنار در آمد و سود از هم مستقل نیستند و استفاده از دوتای آنها کافی است.

۲- با شناخت متغیرها نوع شرکت را بتوان شناسایی کرد. برای مثال با دانستن میزان هزینه‌های تحقیق و توسعه می‌توان فهمید که شرکت چقدر بر دانش فنی اختصاصی تمرکز دارد.

۳- در مدل ما، دیدگاه استراتژیک کل‌نگر در نظر گرفته شده است. در مقابل می‌توان مدل‌های خاص‌تری نیز طراحی نمود که متغیرهای دیگری را بر می‌گزیند. برای مثال اگر تمرکز ما بر گروه‌بندی شرکت‌ها بر اساس ساختارهای مالی آن شرکت بود

1. Hyundai consolidated

2. Input

3. Output

جدول ۳: داده‌های مربوط به ۲۳ شرکت خودروسازی

DMUs	Inputs			Outputs			
	R&D	Total Assets	Average Employees	Production	Value Added	Net Profit / (Loss)	Revenue
1	2,544.0	79,057.0	103.7	1,366.8	12,429	2,874.0	48,999.0
2	5,331.0	190,022.0	365.8	4,589.1	29,706	3,227.0	151,589.0
3	1,401	58,303	170.1	2,173.2	9,260	1,151	51,832
4	317.1	9,117.1	26.6	588.2	215	105.6	9,982.3
5	315.8	6,119.1	12.4	466.7	139	69.5	6,600.1
6	3,450.9	71,479.5	141.3	3,442.0	5,588	4,036.8	66,992.1
7	489.0	10,101.5	33.0	1,140.7	2,246	32.8	14,551.4
8	367.1	5,866.8	7.5	214.0	360	314.2	6,206.3
9	372.6	7,902.0	21.1	621.0	610	398.7	10,695.6
10	647.3	12,093.9	36.2	904.2	873	451.1	19,742.1
11	540.9	9,435.9	19.6	904.0	435	74.3	13,740.0
12	3,026.4	77,630.7	183.5	3,340.8	5,644	3,502.8	63,748.6
13	2,175.0	69,050.0	210.2	3,365.9	9,518	63.0	56,594.0
14	1,963.0	68,766.0	131.0	2,385.1	3,531	2,943.0	41,528.0
15	608.0	12,506.7	40.1	2,200.0	848	445.9	18,569.9
16	82	3,147	29.6	456.3	761	296	4,108
17	323.2	6,945.5	31.3	1,142.5	400	226.7	9,114.2
18	5,494.7	194,266.3	275.9	7,711.0	14,260	9,277.9	142,239.3
19	1,982.0	18,910.0	52.3	926.2	5,514	1,343.0	31,142.0
20	157.9	3,532.7	27.2	556.3	933	371.4	6,838.3
21	4,588.0	136,603.0	328.6	5,659.6	20,065	2,749.0	104,875.0
22	4,588.0	72,085.0	310.6	5,659.6	17,803	1,287.0	96,004.0
23	30.0	2,104.3	30.0	530.2	972	561.4	3,683.6

برای ۹ واحد موثر را نشان می‌دهد و از آنجایی که هیچ کدام از این ۹ واحد با یکدیگر وابستگی خطی ندارند در نتیجه ۹ خوشه خواهیم داشت:

برای دسته‌بندی این واحدها با استفاده از خوشه‌بندی بر اساس DEA، همان‌طور که در بخش قبل توضیح داده شد با استفاده از مدل‌های DEA (CCR) جواب‌های بهینه را به دست می‌آوریم و سپس با استفاده از الگوریتم ارائه شده به خوشه‌بندی واحدهای خودروسازی می‌پردازیم. تابلوی زیر جواب‌های بهینه غیر صفر

جدول ۴: وزن‌های بهینه غیر صفر به دست آمده از DEA

DMUs\Weights	v_1^*	v_2^*	v_3^*	u_1^*	u_2^*	u_3^*	u_4^*
DMU ₁	0.00015928	0.0000036	0.00299344	0.00000991	0.00002746	0.00006284	0.00000948
DMU ₂	0.00008091	0.0000009	0.0010856	0.00000074	0.0000107	0.0000199	0.00000405
DMU ₆	0.00008257	0.00000433	0.00287225	0.00002396	0.00000919	0.00009422	0.0000063
DMU ₈	0.00107017	0.00005607	0.03722835	0.0003106	0.0001191	0.00122126	0.00008168
DMU ₉	0.00092236	0.0000321	0.01911649	0.00010405	0.00014614	0.00037694	0.00006507
DMU ₁₀	0.00049306	0.00000984	0.01553981	0.00007422	0.00008074	0.00016112	0.00004
DMU ₁₁	0.00074276	0.00001483	0.02340986	0.0001118	0.00012163	0.00024272	0.00006026
DMU ₁₈	0.00009932	0.00000077	0.00110686	0.00000514	0.00000477	0.00002281	0.00000479
DMU ₁₉	0.00022633	0.00000628	0.00827204	0.00004034	0.00003187	0.00010391	0.00002079

معنی که این واحدها همان واحدهای تباہیده‌ایی هستند که باید به نزدیک‌ترین دسته مجددا خوشه‌بندی شوند.

واحدهای باقیمانده در واقع همان واحدهایی هستند که حداقل یکی از ضرایب u و یا v در آنها صفر است. به این

جدول ۵: واحدهای تباهیده

DMUs\Weights	v_1^*	v_2^*	v_3^*	u_1^*	u_2^*	u_3^*	u_4^*
DMU ₂₂	0.00007154	0.00000386	0.00126676	0	0.000047	0	0.0000007
DMU ₃	0.00046973	0	0.00201035	0	0.00001062	0	0.0000174
DMU ₁₃	0.00029546	0	0.00170057	0.00001506	0.00001354	0	0.00001052
DMU ₂₁	0.00011688	0	0.00141127	0.00001231	0.0000147	0.00001932	0.00000385
DMU ₁₂	0.00008811	0.00000392	0.00233677	0	0.00000697	0.00008666	0.00000694
DMU ₁₄	0.00012395	0.00000514	0.00307735	0	0	0.00012882	0.00000938
DMU ₄	0.00224027	0	0.0109104	0.00004803	0	0	0.00009123
DMU ₁₅	0	0.00007315	0.0021228	0.00020515	0	0	0.00002955
DMU ₁₆	0.0047747	0	0.02061281	0.00023171	0	0	0.00017341
DMU ₁₇	0	0.00011982	0.0053624	0.00077898	0	0	0
DMU ₂₀	0	0.00028307	0	0	0	0	0.00014624
DMU ₂₃	0	0.00047522	0	0	0.0002038	0	0.0002177
DMU ₅	0.00111734	0	0.05240209	0.00040918	0	0.0005162	0.00008894

در نهایت با استفاده از الگوریتم ارائه شده در بخش قبل به خوشه‌های ارائه شده در جدول زیر می‌رسیم.

جدول ۶: خوشه‌بندی ۲۳ شرکت خودروسازی

Clusters	1	2	3	4	5	6	7
1	DMU ₁	DMU ₂₂					
2	DMU ₂	DMU ₃	DMU ₁₃	DMU ₂₁			
3	DMU ₆						
4	DMU ₈	DMU ₁₂	DMU ₁₄				
5	DMU ₉						
6	DMU ₁₀						
7	DMU ₁₁						
8	DMU ₁₈	DMU ₄	DMU ₁₅	DMU ₁₆	DMU ₁₇	DMU ₂₀	DMU ₂₃
9	DMU ₁₉	DMU ₅					

جدول ۶ قرار می‌گیرند، شرکت‌هایی هستند که خروجی‌هایشان روی فروش بیشتر تمرکز دارد و خوشه ۴ خروجی‌هایش روی سود بیشتر تمرکز دارد. شرکت سایپا به‌عنوان شرکتی که قرار است برایش سیاست بهبود تدوین شود در خوشه هشتم قرار دارد و نسبت به بقیه شرکت‌های این خوشه تمرکز بیشتری روی کل دارایی‌ها دارد. از سوی دیگر بر اساس سیاست‌های این شرکت پیش‌بینی می‌شود که تمرکز شرکت روی تحقیق و توسعه بیشتر خواهد بود و ممکن است وارد خوشه‌ای شود که این ورودی وزن بیشتری دارد (مثلاً خوشه ۴). در این صورت لازم است تمرکز خروجی‌ها از روی فروش به ارزش افزوده متمایل شود. اینها

همانطور که در جدول ۶ نشان داده شده است ۲۳ واحد خودروسازی به ۹ دسته خوشه‌بندی شده‌اند و تابع تولید هر خوشه با توجه به جدول ۵ مشخص است. به عنوان مثال تابع تولید خودروسازی سایپا در خوشه ۸ عبارت است از:

$$0.00000514y_1 + 0.00000477y_2 + 0.0000228y_3 + 0.00000479y_4 + 0.00009932x_1 + 0.00000077x_2 - 0.00110686x_3$$

با نرمال کردن ضرایب توابع تولید می‌توان به مقایسه آنها در خوشه‌ها پرداخت و شباهت توابع تولید در یک خوشه را مشاهده نمود. برای مثال شرکت‌هایی که در خوشه هشتم از خوشه‌بندی

Charnes A, Cooper WW, Golany B, Seiford L, Stutz J. (1985). Foundations of data envelopment analysis for Pareto-Koopmans efficient empirical production functions. *Journal of Econometrics*, 30, 91-107.

Cook WD, Seiford LM. (2009). Data envelopment analysis (DEA) – Thirty years on. *European Journal of Operational Research*, 192, 1-17

Emrouznejad A, Parker BR, Tavares G. (2008). Evaluation of research in efficiency productivity: A survey and analysis of the first 30 years of scholarly literature in DEA. *Socio- Economic planning*, 42(3), 151-157.

Garcia-Palomares UM, Gonzalez-Castano FJ, Burguillo-Rial JC. (2006). A combined global and local search (cgls) approach to global optimization, *Journal of Global Optimization*, 34 (3), 409-426.

Lee C-H et al. (2008). Clustering high dimensional data: A graph-based relaxed optimization approach, *Information Sciences*, 178, PP.4501-4511

Po RW, Guh YY, Yang MS, (2009). A new clustering approach using data envelopment analysis. *European Journal of Operational Research*, 199, 276-284

Sherali HD, Desai J, (2005). A global optimization rlt-based approach for solving the fuzzy clustering prob, *Journal of Global Optimization*, 33 (4), 597-615.

Van der Hofstad R, Kager W. (2008). Pattern theorems, ratio limit theorems and gumbel maximal clusters for random fields, *Journal of Statistical Physics*, 130 (3), 503-522.

Xu R, Wunsch D. (2006). Survey of Clustering Algorithms. *IE Neural Networks*, 16 (3), 645-678

Yin X, Han J, Yu PS, Crossclus. (2007). user-guided multi-relational clustering, *Data mining and Knowledge Discovery*, 15 (3), 321-348.

نمونه‌هایی از تحلیل‌هایی هستند که شرکت سایپا در تدوین سیاست‌های آینده می‌تواند بهره‌برداری نماید.

نتیجه‌گیری

ارتقاء عملکرد واحدهای تصمیم‌گیری با استفاده از روابط بین توابع تولید آنها یکی از مسائل مهم در تصمیم‌گیری مدیران است. در دسته‌بندی بر اساس DEA به دنبال شناسایی روابط پنهان بین عوامل ورودی و خروجی واحدها در تبیین تابع تولید آنها هستیم. یکی از مهم‌ترین نیازهای مدیران به خصوص مدیران واحدهای تولیدی، شناسایی شباهت‌های موجود میان توابع تولید واحدهای مختلف به منظور چگونگی استفاده از منابع در ترکیب ورودی‌ها است. خوشه‌بندی بر اساس تحلیل پوششی داده‌ها با استفاده از توابع به دست آمده از داده‌های ورودی و خروجی می‌تواند برطرف‌کننده این نیاز برای صنایع مختلف و همچنین صنعت خودروسازی باشد. در نتیجه به کارگیری این روش، نه تنها دسته‌ای که هر واحد خودروسازی بدان تعلق دارد مشخص می‌شود، بلکه نوع تابع تولید واحدهای هر دسته نیز شناسایی می‌گردد. از طرف دیگر، با مقایسه تابع تولید به کاررفته در خوشه‌های متفاوت، اطلاعات مهمی در مورد چگونگی استفاده از منابع در ترکیب ورودی‌ها برای صنایع متفاوت به دست می‌آید. با استفاده از خوشه‌بندی بر اساس DEA می‌توانیم خوشه‌های مفیدتر و بامعناتری نسبت به تابع تولید در واحدهای تولیدی داشته باشیم. با استفاده از خوشه‌های به دست آمده از این مدل، می‌توان برای هریک از واحدها یک سیاست بهبود نسبت به چگونگی ترکیب ورودی‌ها و خروجی‌ها تدوین نمود.

Reference

Banker RD, Charnes A, Cooper WW. (1984). Some models for estimating technical and scale inefficiencies in data envelopment analysis. *Management Science* 30, 1078-1092.