

دانشگاه صنعتی شاهرود

دانشکده ریاضی

گروه ریاضی کاربردی

بررسی مدل‌های تلفیقی مکانیابی-تحلیل پوششی داده‌ها

دانشجو: شهزاد رادکانی

استاد راهنما:

دکتر جعفر فتحعلی

استاد مشاور:

دکتر احمد نراکتی رضازاده

پایان نامه ارشد جهت اخذ درجه کارشناسی ارشد ریاضی کاربردی

دی ماه ۱۳۸۸



تقدیم خالصانه به پدر و مادر

و همه آنان که اینچنین اند

مهربان و بی ادعا

قدردانی و تشکر

سپاس خاضعانه به پیشگاه کلام آفرین، که جوانه‌ی واژه‌ها به مدد نغمه‌ی خداوندیش به گل می‌نشیند. سپاس و ستایش پروردگاری را که رحمت بی‌پایانش روشنی بخش دل و دیده‌ام گردید و نعمت وجود خانواده‌ای مهربان را به من ارزانی داشت که در سایه‌ی این نعمت‌ها توانستم از گلستان علم و معرفت اساتید و دانش‌پژوهان توشه‌ای گرفته و از خرمن اطلاعاتشان خوشه‌ای برچینم و گامی گرچه بسیار کوچک در جهت خدمت به دانش دوستان بردارم. اینک که با عنایت و یاری خداوند تدوین و نگارش این پایان‌نامه به پایان رسیده بر خود لازم می‌دانم تا از تمامی کسانی که مرا در این راه یاری نمودند قدردانی نمایم، خصوصاً استاد گرامی آقای دکتر جعفر فتحعلی که در تمام مراحل انجام این پایان‌نامه با راهنمایی‌ها و مساعدت‌های بی‌دریغ و بذل و توجهی مشفقانه پاسخگوی تمامی سوالات من بوده‌اند، و بر خود لازم می‌دانم از زحمات بی‌دریغ آقای دکتر احمد نزاکتی رضازاده کمال تشکر را نمایم که در نهایت صبر و شکیبایی و در تمام مراحل تحصیل در این مقطع و تدوین پایان‌نامه، راهنمایم بوده‌اند و مرا همواره از راهنمایی‌های علمی خود بهره‌مند نموده‌اند. همچنین از آقایان دکتر نادر جعفری راد و دکتر مهدی زعفرانی که قبول زحمت نموده و داوری این پایان‌نامه را به عهده گرفته‌اند تشکر می‌نمایم. و نیز از سایر اساتید گروه ریاضی دانشگاه صنعتی شاهرود که افتخار شاگردی آنها را داشته و یا به نحوی از محضرشان کسب فیض نموده‌ام سپاسگزاری می‌نمایم.

در پایان از خانواده‌ی عزیزم به خصوص پدر و مادر که همیشه و در تمامی مراحل زندگی راهنما و مشوق من بوده‌اند و تمامی موفقیت‌های من مرهون زحمات و فداکاری ایشان می‌باشد، سپاسگزارم. امیدوارم این پایان‌نامه برای اهل فن و دوستداران دانش مفید واقع شود.

چکیده

تاکنون روی مسائل مکانیابی و توسعه و گسترش این دسته مسائل کارهای زیادی انجام شده است همچنین روی مدل‌های اولیه‌ی تحلیل پوششی داده‌ها، در جهت معرفی مدل‌های جدید با توجه به شرایط نو، تحقیقات بسیاری صورت گرفته است. اما تحقیق برای تلفیق این دو مسئله و یافتن مراکز سرویس با بیشترین کارایی و کمترین هزینه کاری جدید است که توسط کلیمبرگ^۱ و راتیج^۲ در سال ۲۰۰۸ انجام شد [۲۲]. تا قبل از آن هیچ مطالعه‌ای در جهت اینکه این دو هدف همزمان برای یافتن و ارتقاء مراکز سرویس بکار رفته باشد، انجام نشده است.

نویسنده‌ی این پایان‌نامه سعی کرده است که در ابتدا به معرفی مدل ارائه شده در مقاله‌ی بالا بپردازد و پایه و اساس پایان‌نامه خود را این مقاله قرار دهد و بعد از تشریح کامل آن، برای بهبود مدل‌های ارائه شده در این مقاله، مدل‌هایی را با شرایط و شکل‌های جدید ارائه دهد. در حقیقت در این پایان‌نامه فرض بر این است که تعدادی نقاط به عنوان سرویس‌دهنده کاندیدای تخصیص به تعدادی سرویس گیرنده می‌باشند. هزینه‌ی استقرار و هزینه‌ی سرویس هر سرویس‌دهنده به هر یک از سرویس‌گیرنده‌ها از پیش تعیین شده است. هر جفت سرویس‌دهنده—سرویس‌گیرنده را به عنوان یک واحد تصمیم‌گیری معرفی کرده و برای ارزیابی کارایی هر واحد شاخص‌هایی معرفی می‌کنیم. هدف از تلفیق مسئله‌ی مکانیابی و تحلیل پوششی داده‌ها دستیابی به مجموعه‌ای از واحدها است که این مجموعه بیشترین مجموع کارایی و کمترین هزینه‌ی سرویس را دارا باشد. ما در این پایان‌نامه با توجه به ارکان موضوع مورد تحقیق در ابتدا به معرفی مسئله‌ی مکانیابی و بطور مشخص دو مدل تخصیص مراکز سرویس با

^۱ Klimberg

^۲ Ratiek

ظرفیت سرویس محدود و نامحدود (UFLP-CFLP) می‌پردازیم و پس از آن با توجه به گستردگی بسیار زیاد مبحث روش تحلیل پوششی داده‌ها سعی می‌کنیم بصورت کلی و جامع به این مبحث بپردازیم و سعی می‌کنیم روش تحلیل پوششی داده‌ها و بخصوص مدل (CCR) به شکلی در این قسمت توضیح داده شود که خواننده پس از مطالعه‌ی این فصل به یک شناخت نسبی از این روش دست یابد. با توجه به اینکه در ادامه روند پایان‌نامه ما نیاز به شکل‌های مختلف مدل‌های تحلیل پوششی داده‌ها داریم در فصل دوم پایان‌نامه به معرفی این مدل‌ها می‌پردازیم و آنها را تشریح می‌کنیم. در فصل سوم به عنوان کاربرد تحلیل پوششی داده‌ها در مسائل مکانیابی الگوریتمی را برای تعیین وزن رؤس در اینگونه مسائل معرفی می‌کنیم که تلفیقی از روش تحلیل پوششی داده‌ها و تحلیل سلسله مراتبی است و آن را برای یک مسئله واقعی پیاده‌سازی می‌کنیم. در فصل چهارم که فصل اصلی پایان‌نامه محسوب می‌شود به معرفی شکل‌های مختلف مدل تلفیقی مسئله‌ی مکانیابی (تخصیص مراکز سرویس) و روش تحلیل پوششی داده‌ها می‌پردازیم.

کلید واژه: مکانیابی؛ تخصیص؛ تحلیل پوششی داده‌ها (DEA)؛ وزن نسبی؛ تحلیل پوششی داده‌های

آینده‌نگر

لیست مقالات مستخرج از پایان نامه

[۱] رادکانی، ش.، فتحعلی، ج.، (۱۳۸۸)، «مدل تلفیقی مکانیابی مراکز سرویس با ظرفیت نامحدود

و روش تحلیل پوششی داده‌ها با در نظر گرفتن وزن نسبی شاخص‌ها»، دومین همایش تخصصی

ریاضی، ساری، دانشگاه پیام نور.

[۲] رادکانی، ش.، فتحعلی، ج.، (۱۳۸۸)، «مدل تلفیقی مکانیابی مراکز سرویس با ظرفیت نامحدود و

روش تحلیل پوششی داده‌ها با در نظر گرفتن ورودی‌ها و خروجی‌های بازه‌ای»، دومین کنفرانس

بین‌المللی تحقیق در عملیات ایران، بابل‌سر، دانشگاه مازندران.

[۳] رادکانی، ش.، حیدری، م.، فتحعلی، ج.، (۱۳۸۸)، «مدل ترکیبی مسئله مکانیابی – تحلیل پوششی

داده‌های تصادفی»، چهارمین کنفرانس ریاضی ایران، تهران، دانشگاه صنعتی شریف.

[۴] جمالیان، ع.، رادکانی، ش.، فتحعلی، ج.، (۱۳۸۸)، «حل مسئله مکانیابی و تخصیص مراکز

سرویس دارای ظرفیت نامحدود با رویکرد ارزیابی کارایی مراکز»، چهارمین کنفرانس ریاضی ایران،

تهران، دانشگاه صنعتی شریف.

فهرست مندرجات

۱	معرفی مسئله‌ی مکانیابی با تأکید بر مسئله‌ی تخصیص	۱
۲	۱-۱ تاریخچه	۲
۴	۲-۱ دسته‌بندی مسائل مکانیابی	۴
۵	۳-۱ اهداف مسائل مکانیابی	۵
۶	۴-۱ اشتباهات متداول در مطالعات مکانیابی	۶
۸	۵-۱ انواع مسائل مکانیابی	۸
۸	۱-۵-۱ مسئله p -میانه (مسئله وبر)	۸
۹	۲-۵-۱ مسئله p -مرکز	۹

۹	۱-۵-۳ مسئله مکانیابی مراکز با ظرفیت سرویس نامحدود (UFLP)
۱۱	۱-۵-۴ مسئله مکانیابی با ظرفیت سرویس محدود (CFLP)
۱۳	۲ تحلیل پوششی داده‌ها
۱۴	۱-۲ مقدمه و تاریخچه
۱۵	۲-۲ تعاریف و مفاهیم اولیه
۱۶	۱-۲-۱ توابع تولید
۱۶	۲-۲-۲ روش پارامتری
۱۷	۳-۲-۲ روش غیر پارامتری
۱۷	۴-۲-۲ بنگاه مرجع
۱۸	۵-۲-۲ مرز کارایی
۱۸	۶-۲-۲ بازده به مقیاس تولید
۱۹	۳-۲ تحلیل پوششی داده‌ها (DEA)
۲۱	۱-۳-۲ اندازه‌گیری کارایی در حالت ورودی محور
۲۲	۲-۳-۲ اندازه‌گیری کارایی در حالت خروجی محور
۲۳	۳-۳-۲ اندازه‌گیری کارایی در تحلیل پوششی داده‌ها
۲۵	۴-۳-۲ وزن‌ها و محدودیت وزنی

۲۸	۵-۳-۲ مثال عددی
۳۱	۴-۲ تحلیل پوششی داده‌های بازه‌ای
۳۵	۵-۲ تحلیل پوششی داده‌های تصادفی
۳۸	۱-۵-۲ مدل تحلیل پوششی داده‌ها محدود شده به قیود تصادفی
۳۸	۲-۵-۲ مدل رضایت بخشی و مفهوم آن در تحلیل پوششی داده‌ها
۳۹	۳-۵-۲ مدل بنکر، چارنزو کوپر (BCC) اصلاح شده تصادفی
۴۲	۴-۵-۲ مدل تحلیل پوششی داده‌های آینده‌نگر
۴۶	۶-۲ محدودیت‌های وزنی در روش تحلیل پوششی داده‌ها
۴۹	۱-۶-۲ اعمال محدودیت‌های کنترل «وزن نسبی» شاخص‌ها در مدل DEA
۵۱	۲-۶-۲ اعمال محدودیت‌های کنترل «وزن نسبی» با تقریب‌های مختلف
۵۴	۳ کاربرد تحلیل پوششی داده‌ها در تعیین وزن رؤس در مسئله‌ی مکانیابی
۵۵	۱-۳ مقدمه
۵۶	۲-۳ فرآیند تحلیل سلسله مراتبی

۵۹	۳-۳	مراحل الگوریتم با رویکرد (AHP/DEA)
۶۰	۳-۳-۱	مرحله ی اول: تشکیل ماتریس مقایسات زوجی با استفاده از DEA .
۶۲	۳-۳-۲	مرحله ی دوم: رتبه بندی با استفاده از AHP.
۶۲	۳-۴	مزایای مدل تلفیقی تحلیل پوششی داده ها و تحلیل سلسله مراتبی
۶۳	۳-۵	مثال عددی
۶۹	۴	مدل تلفیقی مکانیابی-تحلیل پوششی داده ها
۷۰	۴-۱	مدل تلفیقی مکانیابی-تحلیل پوششی داده های قطعی
۷۴	۴-۱-۱	معرفی روش وزن دهی به اهداف برای حل مسائل چندهدفه
۷۶	۴-۱-۲	مثال عددی
۷۹	۴-۲	مدل تلفیقی مکانیابی-تحلیل پوششی داده های بازه ای
۸۱	۴-۳	مدل تلفیقی مکانیابی-تحلیل پوششی داده های تصادفی
۸۳	۴-۳-۱	مثال عددی
۸۵	۴-۴	مدل تلفیقی مسئله مکانیابی-DEA با در نظر گرفتن وزن نسبی شاخص ها . . .

۸۷	۴-۵ یک مدل پیشنهادی برای مسئله‌ی مکانیابی-تحلیل پوششی داده‌ها
۸۹	۴-۵-۱ مثال عددی
۹۱	۵ نتیجه‌گیری، پیشنهادات و کارهای آینده
۹۴	A مراجع
۱۰۳	B منطق فازی
۱۰۸	C واژه نامه
۱۱۳	D فهرست الفبائی

لیست اشکال

۲۲	کارایی در حالت ورودی محور	۱-۲
۲۳	کارایی در حالت خروجی محور	۲-۲
۳۰	مرز کارایی برای مثال عددی	۳-۲
۴۱	مرز کارایی در مدل های DEA و $SDEA$	۴-۲
۵۷	مراحل اجرایی AHP	۱-۳
۵۸	نمایش گرافیکی سلسله مراتب	۲-۳

۳-۳	نمودار شبکه‌ی شهرها	۶۴
۱-۴	نمایش گرافیکی تخصیص با ظرفیت نامحدود با رویکرد هزینه و کارایی	۷۸
۲-۴	نمایش گرافیکی تخصیص با ظرفیت سرویس محدود با رویکرد هزینه و کارایی	۷۸
۳-۴	نمایش گرافیکی تخصیص در حالت تصادفی با رویکرد هزینه و کارایی	۸۴
۴-۴	نمایش گرافیکی تخصیص با رویکرد هزینه و کارایی	۸۹
۱-B	تابع عضویت $\bar{A} = \langle m, \alpha, \beta \rangle$	۱۰۶

لیست جداول

۱.۲	داده‌های ورودی مثال	۲۹
۲.۲	نتایج حاصل از حل مثال	۳۰
۳.۲	مشخصات واحد الگو	۳۱
۴.۲	تفاوت بین مدل‌های تحلیل پوششی داده‌ها و تحلیل پوششی داده‌های تصادفی	۴۰
۱.۳	ماتریس مقایسات زوجی برای تعیین وزن گزینه‌ها برای هر $i = ۱, ۲, ۳$	۵۷
۲.۳	ماتریس مقایسات زوجی برای تعیین وزن معیارها	۵۸

۳.۳	معرفی ورودی‌ها و خروجی‌ها	۶۴
۴.۳	داده‌های ورودی و خروجی برای شهرها	۶۵
۵.۳	ماتریس مقایسات زوجی برای شهرها	۶۸
۶.۳	وزن نهایی گزینه‌ها به روش میانگین حسابی	۶۸
۱.۴	داده‌های مربوط به شاخص‌ها	۷۷
۲.۴	هزینه‌ی استقرار و میزان سرویس مورد نیاز	۷۷
۳.۴	نتایج حل مثال با ظرفیت سرویس نامحدود با در نظر گرفتن $\alpha = 0, \beta = 1$. . .	۷۸
۴.۴	نتایج حل مثال با ظرفیت سرویس نامحدود با در نظر گرفتن $\alpha = 1, \beta = 0$. . .	۷۸
۵.۴	ظرفیت سرویس هر سرویس‌دهنده	۷۸
۶.۴	نتایج حل مثال با ظرفیت سرویس محدود با در نظر گرفتن $\alpha = 0, \beta = 1$. . .	۷۸

۷.۴ نتایج حل مثال با ظرفیت سرویس محدود با در نظر گرفتن $\alpha = 1, \beta = 0$. . . ۷۹

۸.۴ داده‌های مربوط به شاخص‌های تصادفی ۸۴

۹.۴ نتایج حل مثال (۱-۳-۴) با در نظر گرفتن $\alpha = 0, \beta = 1$ ۸۴

۱۰.۴ نتایج حل مثال (۱-۳-۴) با در نظر گرفتن $\alpha = 1, \beta = 0$ ۸۴

۱۱.۴ محاسبه‌ی هزینه و کارایی به شکل مستقل ۹۰

۱۲.۴ نتایج حل مثال (۱-۵-۴) با در نظر گرفتن $\alpha, \beta = 0.5$ ۹۰

فصل ۱

معرفی مسئله‌ی مکانیابی با تأکید بر مسئله‌ی

تخصیص

۱-۱ تاریخچه

مطالعات مکانیابی یکی از اقدام‌های کلیدی در فرایند تعیین محل تسهیلات محسوب می‌شود که توجه به این مهم در موفقیت مراکز، نقش بسزایی دارد. اهمیت این مطالعات به اندازه‌ای است که به تازگی در مورد مراکز فعال نیز، این مطالعات دوباره صورت می‌گیرد و در برخی از موارد منجر به تغییر محل تسهیلات نیز می‌شوند. به طور کلی می‌توان گفت؛ وقتی می‌خواهیم یک سری تجهیزات و تسهیلات را در یک منطقه به گونه‌ای مستقر نمائیم بطوری که بهترین و بیشترین استفاده از این حداقل امکانات را برده باشیم، از علم مکانیابی جهت رسیدن به اهداف خود سود می‌بریم. از قدیم الایام افراد بشر در زندگی روزمره خود بصورت ناخودآگاه از مکانیابی استفاده می‌کردند. برای مثال امپراتور کنستانتین در قرن چهارم قبل از میلاد مسأله مکانیابی نقاطی بر روی شبکه که نشان دهنده مکان پایگاه‌های ارتش روم بود، را حل کرد. پایگاه‌ها باید به گونه‌ای قرار می‌گرفتند که بتوانند در مقابل تهدیدات مهاجمان و شورش‌های محلی از امپراطوری دفاع کنند [۳۱]. نویسنده متذکر می‌شود که در متن جاری منظور از مرکز همان سرویس‌دهنده می‌باشد. به عبارت دیگر مکانیابی مراکز را می‌توان این‌گونه تعریف کرد: انتخاب مکان برای یک یا چند مرکز با در نظر گرفتن سایر مراکز و محدودیت‌های موجود، به گونه‌ای که هدف ویژه‌ای بهینه گردد. این هدف می‌تواند هزینه حمل و نقل، ارائه خدمات عادلانه به مشتریان، در دست گرفتن بزرگترین بازار و غیره باشد [۵۱]. مسأله مکانیابی توسط فرما^۱ در قرن هفدهم بصورت زیر مطرح شد: فرض کنید سه نقطه در صفحه داده شده است، نقطه چهارم را بگونه‌ای بیابید که مجموع فاصله‌های آن تا سه نقطه داده شده، کمینه شود. توریچلی^۲ در سال ۱۶۴۰ این مسئله را حل کرد. به این دلیل نقطه‌ی بهینه را نقطه‌ی توریچلی و مسئله را مسئله فرما نامیدند.

^۱ Fermat^۲ Torricelli

در سال ۱۹۰۹ اولین تعریف مسئله مکانیابی بصورت کاربردی و مدرن توسط وبر^۳ ارائه شد [۴۰]. در سال ۱۹۶۴ حکیمی مطالعات جدی را بر روی این مسئله انجام داد و تابع هدف این مسائل را به صورت کمترین مجموع^۴ و کمترین بیشینه^۵ مطرح کرد او همچنین مسائل مکانیابی را بر روی شبکه‌ها نیز بررسی کرد [۱۷، ۱۶]. اولین طبقه بندی مدل‌های مختلف مکانیابی توسط هندلر^۶ و میرچندانی^۷ ارائه شد [۱۵].

همچنین طبقه بندی های دیگری در این زمینه توسط کراروپ^۸ و پروزن^۹ [۲۴]، هانسن^{۱۰} و همکاران [۱۹]، میرچندانی و فرانسیس^{۱۱} [۲۶]، فرانسیس و همکاران [۱۴]، اُون^{۱۲} و دسکین^{۱۳} [۲۸]، اسکاپارا^{۱۴} و اسکاتالا^{۱۵} [۳۵]، درزنر^{۱۶} و هاماختر^{۱۷} [۱۰] انجام شده است.

Weber^۳Minisum^۴Minimax^۵Handler^۶Mirchandani^۷Krarup^۸Pruzan^۹Hansen^{۱۰}Francis^{۱۱}Owen^{۱۲}Daskin^{۱۳}Scaparra^{۱۴}Scutella^{۱۵}Drezner^{۱۶}Hamacher^{۱۷}

۱-۲ دسته‌بندی مسائل مکانیابی

مسائل و مدل‌های مکانیابی را می‌توان به روش‌های متفاوتی طبقه‌بندی نمود. انواع گوناگونی از مسائل مکانیابی تسهیلات وجود دارند که برخی از طبقه‌بندی‌های مهم آن عبارتند از:

(۱) مسئله مکانیابی تسهیلات گسسته؛ که در آن مجموعه‌ای از نقاط تقاضا و مکان‌های تسهیلات مشخص می‌باشد.

(۲) مسئله مکانیابی تسهیلات پیوسته؛ که در فضایی مشخص تعداد معینی از تسهیلات را مکانیابی می‌کنند.

(۳) مسئله مکانیابی تسهیلات شبکه‌ای؛ که در یک تعداد معینی از اتصالات و گره‌ها محدود شده‌اند.

(۴) مسئله مکانیابی تسهیلات تصادفی؛ که در آن برخی از پارامترهای تقاضا یا زمان سفر نامعین هستند.

طبقه‌بندی‌ها براساس مشخصه‌های گوناگونی صورت می‌پذیرند. برخی از این مشخصه‌ها عبارتند از:

- درختی یا شبکه‌ای بودن مدل
- تعداد مراکز استقرار
- ساکن یا پویا بودن مدل
- قطعی یا احتمالی بودن مدل
- محدود بودن یا نبودن ظرفیت مراکز استقرار یافته

برخی از عوامل مؤثر بر طبقه‌بندی نیز عبارتند از:

✓ مشخصات تسهیلات جدید: تعداد، اندازه تسهیلات، نقطه‌ای و یا ناحیه‌ای بودن مکان تسهیلات.

✓ مکان تسهیلات فعلی

✓ اندازه‌گیری فاصله‌ها: خط مستقیم، اقلیدسی، پله‌ای.

✓ فضای جواب: تک‌بعدی، چند بعدی، پیوسته، گسسته، محدود، نامحدود.

✓ تأثیر متقابل تسهیلات جدید: رابطه‌ای کیفی / کمی، پارامتری / متغیری، قطعی / احتمالی و ...

✓ نوع هدف: مینیمم/ماکسیمم، کمی/کیفی.

اما در حالت کلی مسائل مکانیابی با توجه به هدف‌های متعدد به دسته‌های گوناگونی تقسیم می‌شوند که در ادامه به بعضی از آنها اشاره می‌کنیم.

۳-۱ اهداف مسائل مکانیابی

مسائل مکانیابی، هدف‌های مختلفی را در بر دارند. هدف‌ها در شناسایی و اولویت بندی معیارهای تصمیم‌گیری در یک مسأله مکانیابی و زیر معیارهای آن‌ها، اهمیت و نقش مهمی دارند. در یک تقسیم‌بندی هدف‌های مسائل مکانیابی با دو رویکرد برنامه‌ریزی ریاضی و بر حسب انواع تابع هدف به سه دسته تقسیم شده‌اند [۱۰]:

(۱) هدف‌های کششی^{۱۸}: این هدف‌ها اشاره به نزدیکی هر چه بیشتر محل استقرار سرویس‌دهنده (تسهیلات) به مشتریان و کمتر کردن مسافت دارند که شامل قدیمی‌ترین مسائل مکانیابی می‌شوند. در واقع مسائلی که تابع هدف آنها بصورت کمینه‌سازی است، هدف‌های کششی دارند.

^{۱۸} Pull objective

۲) هدف‌های فشاری^{۱۹}: این هدف‌ها مسائل مکانیابی مراکز نامطلوب^{۲۰} را در بر می‌گیرند و از اوایل دهه ۱۹۷۰ بوجود آمدند. هدف در این مسائل، حداکثر کردن فاصله مراکز جدید از مراکز موجود است. مدل‌هایی که برای این نوع هدف‌ها ارائه شدند بعدها به مدل‌های مکانیابی ناخوشایند^{۲۱} معروف شدند. مثال برای این هدف‌ها، یافتن مکان مناسب برای دفن زباله است که در آن، یکی از هدف‌ها بهینه کردن فاصله این مکان از مناطق مسکونی است.

۳) هدف‌های متعادل^{۲۲}: هدف‌هایی هستند که تلاش در متعادل ساختن مسافت بین مراکز و مشتریان دارند و هدف اصلی آنها دست‌یابی به برابری است. این هدف‌ها بیشتر در تصمیم‌گیری‌های عمومی کاربرد دارند. جایی که هدف برقراری عدالت بین افراد است. مانند متعادل کردن حجم کاری مراکز پلیس که سبب متعادل شدن ارائه خدمات به متقاضیان می‌شود.

۱-۴ اشتباهات متداول در مطالعات مکانیابی

اشتباه در تعیین محل، ضررهای جبران‌ناپذیری به دنبال خواهد داشت که گاهی منجر به تغییر محل سرویس‌دهنده با صرف هزینه‌های زیاد شده، یا به رکود و تعطیلی سرویس‌دهنده (مثلاً کارخانه) می‌انجامد. عموماً اشتباه در تعیین محل، هنگامی پیش می‌آید که تعریف درستی از آن چه که از ما خواسته می‌شود در دست نباشد. ولی اشتباه‌های دیگری نیز وجود دارد که حتی مدیران زیرک نیز دچار آن می‌شوند. برخی از این نوع اشتباه‌ها برای توجه بیشتر مدیران، محققان و افراد کلیدی و تصمیم‌گیری در مسائل

^{۱۹} Push objective

^{۲۰} Undesirable locaton

^{۲۱} Obnoxious locaton

^{۲۲} Balancing objective

مکانیابی به این شرح بیان می‌شود [۵۵]:

- (۱) فقدان بازرسی و شرح دقیق عوامل و نیازمندی‌ها.
- (۲) چشم‌پوشی از بعضی شرایط مورد نیاز و بررسی ناقص نیازمندی‌های طرح.
- (۳) علایق شخصی یا تعصبات مسئولان در پذیرش حقایق و دلایل منطقی و علمی.
- (۴) مقاومت مدیران اجرایی در انتقال به محل جدید.
- (۵) توجه بیش از اندازه به نواحی شلوغ و صنعتی و در نتیجه نادیده گرفتن ناحیه‌هایی که به تازگی صنعتی شده و یا در شرف صنعتی شدن قرار دارند.
- (۶) توجه بیش از اندازه به هزینه‌های زمین و در نتیجه انتخاب زمین‌های ارزان.
- (۷) بی‌توجهی به هزینه حمل و نقل و عدم برآورد درست آن.
- (۸) قضاوت در مورد نیروی انسانی بالقوه بر مبنای نرخ دستمزد و بدون توجه به کارائی، مهارت، سابقه و تاریخچه کارگری و سایر عوامل مؤثر در انتخاب نیروی انسانی.
- (۹) انتخاب جامعه‌ای با سطح فرهنگ و تحصیلات پایین به گونه‌ای که جذب نیروی متخصص بسیار مشکل باشد.
- (۱۰) پافشاری در منافع آنی و کوتاه مدت و بی‌توجهی به آن.
- (۱۱) کافی نبودن اطلاعات و یا نادرست بودن آن‌ها در مورد بازار، شیوه‌های حمل و نقل، مواد خام و سایر عوامل که در برآورد هزینه‌ها تأثیر دارند.
- (۱۲) عوامل محیطی از جمله فشارهای سیاسی.

- (۱۳) خطا در به کارگیری روش‌ها و تکنیک‌های تصمیم‌گیری مکانیابی.
- (۱۴) عدم اولویت‌بندی (وزن دهی) مناسب به معیارهای تصمیم‌گیری.
- (۱۵) نبود اطلاعات دقیق و کلی در زمینه معیارهای مورد نظر.
- (۱۶) بی توجهی به تغییر و تحولات آینده (تهدیدها، فرصت‌ها، رشد تقاضا، به هم خوردن توازن مناطق).
- بنابراین انجام مطالعات مکانیابی درست و مناسب، علاوه بر تأثیر اقتصادی اثرات اجتماعی، محیط زیستی و فرهنگی را باعث خواهد شد.

۱-۵ انواع مسائل مکانیابی

مسائل مکانیابی دارای تنوع بسیار زیادی هستند. از این رو برای سهولت در بیان، این مسائل را به راه‌های مختلفی دسته‌بندی کرده‌اند، اما به طور کلی مسائل تحلیل مکان در یکی از دسته‌های زیر قرار می‌گیرند:

۱-۵-۱ مسئله p -میان (مسئله وبر)

این قبیل مسائل برای مکانیابی p سرویس‌دهنده، در p مکان انجام می‌شود و یک معیار هزینه‌ای را مینیمم می‌کند. هزینه ممکن است بر حسب زمان، پول، تعداد سفر، مسافت کل یا هر مقیاس دیگری بیان شود. به علت اینکه در این گونه مسائل، هدف حداقل کردن هزینه کل است، با نام مسائل حداقل مجموع یا مسئله‌ی وبر نیز مطرح می‌شوند [۱۶، ۱۸].

۱-۵-۲ مسئله p -مرکز

این مسائل برای تعیین مکان p مرکز به منظور حداقل کردن حداکثر فاصله هر مرکز، تا نقطه تقاضایی که برای خدمت دادن به آن نقطه مورد تقاضا تعیین شده است، استفاده می‌شوند. در واقع این گونه مسائل برای استقرار خدمات اورژانس، مانند: آتش‌نشانی، خدمات آمبولانس و مراکز پلیس در جامعه مورد استفاده قرار می‌گیرند. در این مسائل تعداد مراکز از پیش مشخص است. این مسائل به دو حالت تقسیم می‌شوند. حالت رأسی مسئله‌ی p -مرکز، مسئله را به استقرار تسهیلات جدید در گره‌های شبکه محدود می‌نماید. در حالی که در مدل مطلق مکانیابی p -مرکز، تسهیلات جدید در هر نقطه از شبکه می‌توانند واقع شوند [۱۴].

۱-۵-۳ مسئله مکانیابی مراکز با ظرفیت سرویس نامحدود (UFLP)

مسائل مکانیابی مراکز با ظرفیت سرویس نامحدود ($UFLP^{23}$) در دسته مسائل حداقل مجموع قرار می‌گیرند اما در این مسائل هزینه، هزینه ثابت را نیز شامل می‌شود و هزینه ثابت به مکانی بستگی دارد که مرکز در آن قرار می‌گیرد. تعداد مراکزی که باید استقرار یابند از پیش مشخص شده نیست، اما به گونه‌ای معین می‌شوند که هزینه را کمینه کنند. به علت اینکه در این گونه مسائل ظرفیت هر مرکز نامحدود در نظر گرفته می‌شود، تخصیص کل تقاضا به بیش از یک نقطه تأمین هرگز سود بخش نخواهد بود. مدل

²³ Uncapacitated Facility Location Problem

UFLP به شکل زیر می باشد:

$$\begin{aligned} \min \quad & \sum_{k=1}^K \sum_{l=1}^L C_{kl} \text{dem}_l x_{kl} + \sum_{k=1}^K F_k y_k \\ \text{s.t} \quad & \sum_{k=1}^K x_{kl} = 1 \quad l = 1, \dots, L \\ & x_{kl} \leq y_k \quad l = 1, \dots, L \quad \text{and} \quad k = 1, \dots, K \\ & x_{kl}, y_k \in \{0, 1\} \end{aligned} \quad (1-1)$$

که در آن k اندیس مربوط به سرویس دهنده، l اندیس مربوط به سرویس گیرنده‌ها و C_{kl} هزینه‌ی هر واحد سرویسی است که سرویس دهنده‌ی k به سرویس گیرنده l می‌دهد. dem_l میزان سرویس مورد نیاز سرویس گیرنده‌ی l است. F_k هزینه ثابت (استقرار) سرویس دهنده‌ی k است. x_{kl}, y_k هر دو متغیرهای تصمیم این مسئله هستند که به صورت زیر تعریف می‌شوند:

$$x_{kl} = \begin{cases} 1 & \text{اگر سرویس دهنده‌ی } k \text{ به سرویس گیرنده‌ی } l \text{ سرویس دهد} \\ 0 & \text{در غیر این صورت} \end{cases}$$

$$y_k = \begin{cases} 1 & \text{اگر سرویس دهنده } k \text{ استقرار یابد} \\ 0 & \text{در غیر این صورت} \end{cases}$$

محدودیت اول از نامحدود بودن ظرفیت سرویس سرویس دهنده‌ها نتیجه می‌شود یعنی هر سرویس گیرنده در صورت ارتباط با یک سرویس دهنده تمام نیاز خود را بر طرف می‌سازد. محدودیت دوم نشان دهنده‌ی این موضوع است که هر واحدی برای سرویس دهی باید قبل از آن استقرار یابد و در صورتی که واحدی استقرار نیافته باشد نمی‌تواند عمل سرویس دهی را انجام دهد. تابع هدف در این مسئله به صورت مینیمم سازی هزینه است [۲۳، ۲۶].

۱-۵-۴ مسئله مکانیابی با ظرفیت سرویس محدود (CFLP)

مسائل مکانیابی با ظرفیت سرویس محدود CFLP^{۲۴} شبیه به مسائل UFLP هستند. فقط در این مسائل ظرفیت هر کدام از مراکز محدود است. ممکن است در این مورد جواب بهینه به گونه‌ای باشد که یک مشتری به بیش از یک منبع تأمین، ارجاع داده شود. در واقع ممکن است پس از تخصیص مشتری به یک مرکز، پس از برآوردن بخشی از تقاضای مشتری ظرفیت مرکز به پایان برسد و برای برآوردن باقی مانده تقاضای مشتری مجبور به اختصاص آن به دیگر مراکز که هزینه بیشتری نیز در بر دارند، شویم. البته گاهی ممکن است با وجود اینکه اختصاص یک مشتری به یک مرکز ویژه هزینه کمتری را در بر داشته باشد، به دلیل اینکه ظرفیت آن مرکز توسط مشتریان دیگر پر شده است مجبور به اختصاص کل تقاضای آن مشتری به مراکز دیگر شویم. مدل CFLP به شکل زیر می‌باشد:

$$\begin{aligned}
 \min \quad & \sum_{k=1}^K \sum_{l=1}^L C_{kl} b_{kl} + \sum_{k=1}^K F_k y_k \\
 s.t. \quad & \sum_{k=1}^K x_{kl} \geq 1 & l = 1, \dots, L \\
 & x_{kl} \leq y_k & l = 1, \dots, L \quad \& \quad k = 1, \dots, K \\
 & \sum_{k=1}^K b_{kl} = dem_l & l = 1, \dots, L \\
 & b_{kl} \leq \min[dem_l, Cap_k] y_k & l = 1, \dots, L \quad \& \quad k = 1, \dots, K \\
 & x_{kl}, y_k = \{0, 1\} \\
 & b_{kl} \geq 0
 \end{aligned} \tag{۲-۱}$$

در مدل (۲-۱) پارامتر Cap_k بیان کننده میزان ظرفیت سرویس، سرویس دهنده k می‌باشد و متغیر b_{kl} میزان سرویسی است که سرویس دهنده k به سرویس گیرنده l می‌دهد. توجه داشته باشید همانطور که در محدودیت سوم مشخص است مجموع میزان سرویسی که سرویس دهنده‌های مختلف به یک سرویس گیرنده مشخص می‌دهند نباید از نیاز آن سرویس گیرنده بیشتر باشد. محدود بودن ظرفیت

سرویس، سرویس‌دهنده‌ها منجر به این موضوع می‌گردد که بعضی از سرویس‌گیرنده‌ها تمام نیاز خود را از یک سرویس‌دهنده برطرف نسازند که این موضوع در محدودیت اول نشان داده شده است. محدودیت چهارم این مدل بیان می‌دارد که میزان سرویسی که سرویس‌دهنده به سرویس‌گیرنده می‌دهد باید متناسب با ظرفیت سرویس سرویس‌دهنده و میزان سرویس مورد نیاز سرویس‌گیرنده باشد [۲۲، ۲۳].

فصل ۲

تحلیل پوششی داده‌ها

۱-۲ مقدمه و تاریخچه

تحلیل پوششی داده‌ها^۱ (DEA) یکی از ابزارهای قدرتمند مدیریتی است. تحلیل پوششی داده‌ها با روش قدرتمندی که در دست دارد، قادر است مدیریت را در جهت نیل به اهداف عالی سازمان و در جهت استفاده بهینه از منابع و تخصیص آنها و در نهایت کسب سودآوری بیشتر، یاری رساند. DEA ابزاری در اختیار مدیران قرار می‌دهد، تا بتواند بوسیله آن عملکرد شرکت خود را در قبال سایر رقبا محک زنند، و بر اساس نتایج آن برای آینده‌ای بهتر تصمیم‌گیری کنند.

تاکنون مطالعات و تحقیقات زیادی در انجمن‌های مختلف و دانشگاه‌های مختلف جهان در مورد تحلیل پوششی داده‌ها و کاربردهای آن صورت گرفته است. سادگی فهم و اجرای روش DEA و در کنار دقت بالا و کاربرد وسیع آن در زمینه‌های مختلف سیاسی، فرهنگی، اجتماعی و اقتصادی باعث شده است محققان زیادی از این روش برای دست یافتن به اهداف خود استفاده کنند. آمارهایی که در ذیل به آنها اشاره شده است، خود گواهی بر صحت این سخن است. تا سال ۲۰۰۱ بیش از ۱۷۷ کتاب، ۱۴۶۹ مقاله، ۱۲۷۱ روزنامه و مجله، ۲۸۶ تحقیق و پژوهش در مورد DEA منتشر شده است. در سال‌های اخیر تلاش‌های زیادی برای رفع برخی نواقص روش‌های DEA صورت گرفته است. از جمله این تلاش‌ها، استفاده از منطق فازی؛ تئوری احتمالات و ... برای داده‌های ناقص و کیفی بوده است [۴۶].

تاریخچه‌ی روش تحلیل پوششی داده‌ها به موضوع رساله‌ی رودز^۲ به راهنمایی پروفیسور کوپر^۳ برمی‌گردد که عملکرد مدارس دولتی آمریکا را مورد ارزیابی قرار داد. این مطالعه منجر به چاپ اولین مقاله درباره‌ی معرفی عمومی تحلیل پوششی داده‌ها در سال ۱۹۷۸ میلادی گردید [۳]. در این مقاله سه

^۱ Data Envelopment Analysis

^۲ Rhodes

^۳ Cooper

متخصص تحقیق در عملیات از طریق برنامه‌ریزی ریاضی، اندازه‌گیری کارایی را معرفی کردند روش تحلیل پوششی داده‌ها با جامعیت بخشیدن به روش فارل^۴، بگونه‌ای که خصوصیت فرایند تولید با چند عامل تولید (ورودی) و چند محصول (خروجی) را دربرگیرد به ادبیات اقتصادی اضافه گردید. با پیشرفت و تکامل این روش در حال حاضر تحلیل پوششی داده‌ها یکی از حوزه‌های فعال تحقیقاتی در اندازه‌گیری کارایی بوده و بطور چشمگیری مورد استقبال پژوهشگران جهان قرار گرفته است. این روش برای ارزیابی عملکرد سازمان‌های دولتی و غیرانتفاعی که اطلاعات قیمتی آنها معمولاً در دسترس نیست یا غیر قابل اتکا است کاربرد قابل ملاحظه‌ای دارد [۵۲].

۲-۲ تعاریف و مفاهیم اولیه

تعریف ۱.۲ [۴۶] نسبت خروجی واقعی به خروجی مورد انتظار با مقیاس ورودی واقعی را کارایی می‌گویند.

تعریف ۲.۲ ورودی عاملی است که با افزایش آن، با حفظ تمام عوامل دیگر، کارایی کاهش یافته و با کاهش آن با حفظ تمام عوامل دیگر کارایی افزایش می‌یابد. در واقع رابطه معکوس بین میزان ورودی‌ها و کارایی وجود دارد [۴۶].

تعریف ۳.۲ خروجی عاملی است که با افزایش آن، با حفظ تمام عوامل دیگر، کارایی افزایش یافته و با کاهش آن با حفظ تمام عوامل دیگر، کارایی کاهش می‌یابد. این بدین مفهوم است که رابطه مستقیم

Farrell^۴

بین میزان خروجی‌ها و کارایی وجود دارد [۴۶].

۱-۲-۲ توابع تولید

شناخت تابع تولید یکی از موضوعات مهم علم اقتصاد می‌باشد. تابع تولید تابعی است که بیشترین خروجی ممکن را از ترکیب ورودی‌ها فراهم می‌کند. بنابراین اگر مقدار خروجی را با Q و ورودی‌ها را با $x_1, x_2, \dots, x_n$ نشان دهیم، تابع تولید را می‌توان بصورت $Q = f(x_1, x_2, \dots, x_n)$ در نظر گرفت. در واقع تابع تولید معرف رابطه مقداری موجود بین میزان تولید و عوامل مورد نیاز برای تولید است. از مشکلات عمده این تعریف برای تابع تولید قابل استفاده بودن برای سیستم‌هایی تنها با یک خروجی و دشواری تشخیص ضابطه‌ی f می‌باشد. تشخیص تابع تولید به دو صورت پارامتری و غیر پارامتری صورت می‌گیرد که در ادامه بطور خلاصه در مورد مفاهیم هر یک بحث می‌کنیم [۴۶].

۲-۲-۲ روش پارامتری

در این روش برای تابع تولید پیش فرض اولیه‌ای در نظر گرفته شده، سپس با استفاده از اطلاعات موجود پارامترهای تابع تخمین زده می‌شود. تحلیل رگرسیون یک روش پارامتری است. در تخمین پارامترهای این تابع وقتی بیش از دو عامل متغیر در تابع تولید دخالت دارد استفاده از تکنیک‌های کمی از قبیل روش‌های تحقیق در عملیات اجتناب ناپذیر است. عمده اشکالات روش‌های پارامتری که باعث می‌شود تا اینگونه روش‌ها برای ارزیابی واحدها مناسب نباشد، بقرار زیر است:

(۱) لزوم پیش فرض اولیه برای تابع تولید، ممکن است با ماهیت واحدهای تحت ارزیابی در تضاد باشد.

(۲) در این روش ارزیابی واحدهایی با یک خروجی امکان پذیر است.

۳) ماهیت غیر خطی بودن تابع تولید در این روش‌ها، محاسبات را پیچیده می‌کند [۴۶].

۳-۲-۲ روش غیر پارامتری

این روش با استفاده از داده‌های در دسترس، تابع تولید را تخمین می‌زنند و نیازی به در نظر گرفتن فرضیات محدودیت‌های تئوریک روی تابع تولید نیست. فارل در سال ۱۹۵۷ برای نخستین بار روش غیر پارامتری را مطرح کرد [۱۲]. او با استفاده از خروجی و ورودی‌های واحدهای تصمیم‌گیری^۵ تابع تولید را چنان بر مجموعه‌ی خروجی و ورودی‌ها پردازش داد که حاصل پردازش فوق یک تابع قطعه قطعه خطی بود. مقاله فارل اساس کار مقاله‌ی چارنز^۶، کوپر و رودز در سال ۱۹۷۸ قرار گرفت [۳]. آنها تحلیل اولیه فارل را که در حالت یک خروجی و چند ورودی مطرح شد، به حالت چند ورودی و چند خروجی تعمیم دادند. ادامه دهندگان این روش، بنکر^۷، چارلز و کوپر بودند که در سال ۱۹۸۴ مدل BCC یکی از مدل‌های DEA را ارائه دادند [۱]. این افراد پایه‌گذاران مدل‌های تحلیل پوششی داده‌ها هستند، که در مباحث بعدی با آنها و نیز مدل‌های DEA بیشتر آشنا خواهیم شد.

۴-۲-۲ بنگاه مرجع

یکی از مزیت‌های روش تحلیل پوششی داده‌ها معرفی بنگاه مرجع (مجازی) برای کارا کردن یک بنگاه ناکارآمد است در این روش برای هر بنگاه ناکارآمد، ترکیبی از بنگاه‌های کارآمد، بنگاهی را می‌سازند که الگوی بنگاه مورد نظر قرار می‌گیرد (بنگاه‌های کارا هر کدام با وزن خاصی که از حل مدل بدست می‌آید، در بنگاه مرجع بنگاه ناکارا حضور دارند). با توجه به بنگاه مرجع بدست آمده، مقادیر پیشنهادی

^۵ Decision Making Unit (DMU)

^۶ Charnes

^۷ Banker

برای کارا کردن بنگاه ناکارا بدست می‌آید. [۴۶]

۵-۲-۲ مرز کارایی

مرز کارایی متشکل از واحدهائی با اندازه کارایی یک است که برخی از واحدهای آن واقعی و در حقیقت تجربه شده‌اند و برخی از آنها مجازی هستند به این معنا که هرچند واحد مذکور عینیت نیافته است از اینرو با مجموعه‌ای از واحدهای تجربه شده امکان تحقق چنین واحدهایی وجود دارد که به آنها واحدهای مجازی می‌گویند [۴۶].

۶-۲-۲ بازده به مقیاس تولید

در تئوری اقتصاد خرد بازده به مقیاس عبارت است از تأثیر عوامل تولید بر تولید. فرض کنید تابع تولید بصورت $Q = F(x_1, x_2, \dots, x_n)$ باشد اگر عوامل تولید یعنی همه‌ی ورودی‌ها λ برابر شود و مقدار تابع نیز به اندازه‌ی λ تغییر کند، در این صورت بازده به مقیاس ثابت است به عبارت دیگر

$$\forall \lambda \geq 0 \quad F(\lambda x_1, \lambda x_2, \dots, \lambda x_n) = \lambda F(x_1, x_2, \dots, x_n)$$

اگر افزایش تولید بیشتر از افزایش ورودی‌ها باشد یعنی

$$\forall \lambda > 1 \quad F(\lambda x_1, \lambda x_2, \dots, \lambda x_n) > \lambda F(x_1, x_2, \dots, x_n)$$

$$\forall \lambda < 1 \quad F(\lambda x_1, \lambda x_2, \dots, \lambda x_n) < \lambda F(x_1, x_2, \dots, x_n)$$

گوییم بازده به مقیاس صعودی است و اگر افزایش تولید کمتر از افزایش ورودی‌ها باشد یعنی

$$\forall \lambda > 1 \quad F(\lambda x_1, \lambda x_2, \dots, \lambda x_n) < \lambda F(x_1, x_2, \dots, x_n)$$

$$\forall \lambda < 1 \quad F(\lambda x_1, \lambda x_2, \dots, \lambda x_n) > \lambda F(x_1, x_2, \dots, x_n)$$

گوییم بازده به مقیاس نزولی است. البته تعریف فوق برای حالتی است که چند ورودی جهت تولید

یک خروجی بکار برده شود [۴۶].

۳-۲ تحلیل پوششی داده‌ها (DEA)

تحلیل پوششی داده‌ها یک ابزار کمی، استاندارد و با کاربرد گسترده در مطالعات اندازه‌گیری کارایی و تحلیل عملکرد می‌باشد. DEA، کارایی نسبی واحدهایی که دارای ورودی‌ها و خروجی‌های مشابه می‌باشند را اندازه‌گیری می‌کند. ما اینگونه واحدها را واحدهای تصمیم‌گیرنده (^ADMU) می‌نامیم. DEA کارایی یک DMU را در مقایسه با سایر واحدها ارزیابی می‌کند. به همین خاطر امتیاز کارایی DMU، یک امتیاز نسبی خواهد بود. ارزیابی و مقایسه عملکرد و کارایی واحدهای مشابه قسمت مهمی از مدیریت یک سازمان پیچیده می‌باشد. DEA راه‌کارهایی را برای مدیریت بهتر منابع جهت نیل به خروجی‌های مورد انتظار، ارائه می‌دهد. کارا بودن یا غیر کارا بودن یک DMU بستگی به عملکرد آن واحد در انتقال ورودی‌ها به خروجی‌هایش در مقایسه با سایر واحدها در یک حوزه خاص دارد [۲۹]. DEA اخیراً در مجموعه علوم مدیریت قرار گرفته است. مهمترین علت موفقیت DEA بعنوان یک ابزار کمی، غیر پارامتری بودن روش آن است. هر DMU با استفاده از تعاریف تئوری استاندارد برای محاسبه کارایی، امتیازدهی می‌شود، که این امتیاز بوسیله مقیاس‌های خاص که سعی در حداکثر نمودن امتیاز کارایی آن واحد دارند، محاسبه می‌شود. از ویژگی‌های مهم DEA می‌توان موارد زیر را نام برد:

(۱) قابلیت استفاده و فهم آسان و راحت

(۲) ارزیابی واقع بینانه و ارزیابی توأم مجموعه عوامل دخیل در مدل‌ها

(۳) عدم نیاز به وزن‌ها و تابع تولید از پیش تعیین شده (غیر پارامتری)

(۴) تصویر کردن بهترین وضعیت عملکردی بجای وضعیت مطلوب

(۵) قابلیت وارد نمودن چندین ورودی و چندین خروجی در مدل

قابلیت های کاربردی روش DEA عبارتند از :

- رتبه بندی واحدهای تصمیم گیرنده
- ارائه واحدهایی جهت مقایسه کارایی و ارائه راهکارهای بهبود کارایی
- تعیین پیشرفت و پسرفت تکنیکی واحدها
- تخصیص بهینه منابع
- تحلیل حساسیت ورودی ها و خروجی ها
- تعیین پتانسیل های عملکردی

اما مانند سایر روش ها، DEA نیز دارای معایبی نیز می باشد که می توان برخی از مهمترین آنها را اینگونه بیان کرد:

(۱) در نظر نگرفتن اختلالات تصادفی و عوامل محیطی. که برای رفع این مشکل امروزه محققان سعی کرده اند تا با دخالت دادن اصول احتمال و مفاهیم منطق فازی این مشکل روش DEA را برطرف کنند.

(۲) غیرپارامتری بودن روش DEA که در برخی مواقع منجر به نتایج نه چندان مناسب می شود.

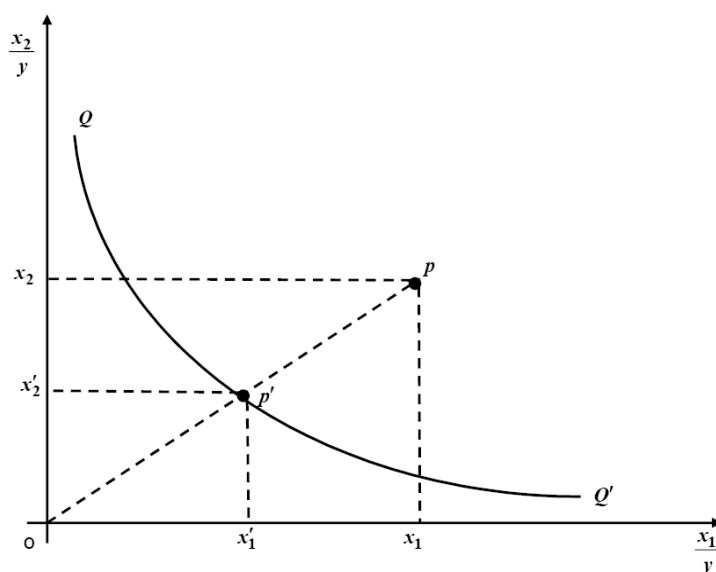
(۳) حساسیت نتایج DEA نسبت به پراکندگی داده ها. پراکندگی داده ها منجر به ارائه نتایج نامطلوب می شود. به همین منظور فرض می شود که پراکندگی داده ها قابل چشم پوشی است.

روش DEA بعنوان روش اندازه‌گیری کارایی تکنیکی^۹ شناخته می‌شود. کارایی تکنیکی سعی دارد که حداکثر خروجی را با مصرف ورودی‌هایی که به واحد داده شده است، ایجاد کند.

۲-۳-۱ اندازه‌گیری کارایی در حالت ورودی محور

مفهوم ورودی محور این است که به چه میزان باید ورودی‌ها را با ثابت نگه داشتن میزان خروجی‌ها کاهش داد تا واحد مورد نظر به مرز کارایی برسد. برای درک بهتر مفهوم اندازه‌گیری ورودی محور فرض می‌کنیم که واحدهایی در اختیار داریم که دو نوع ورودی x_1 و x_2 را برای ایجاد یک خروجی y مصرف می‌کنند مقدار خروجی را یک در نظر می‌گیریم شکل (۲-۱) این مفهوم را نشان می‌دهد. منحنی QQ' مکان هندسی تمام ترکیبات ممکن از ورودی‌ها است که خروجی یکسان تولید می‌کنند. این منحنی نسبت به مبدأ محدب می‌باشد که نشان دهنده‌ی محدوده‌ی پایین ورودی‌ها است و بگونه‌ای است که بتواند تمام واحدها (نقاط داده‌ای) را دربرگیرد. یعنی حداکثر سطح ممکن برای کاهش ورودی‌ها را نشان می‌دهد. بنابراین واحدهایی که روی این منحنی قرار می‌گیرند، واحدهای کارا می‌باشند. زیرا با حداقل ممکن مصرف ورودی به خروجی y دست یافته‌اند. همان‌طور که بیان شد می‌توان نتیجه گرفت که واحد p غیرکارا است. واحد p برای کارا شدن نیاز دارد که میزان ورودی x_1 خود را به میزان ورودی x'_1 و میزان ورودی x_2 خود را به میزان ورودی x'_2 کاهش دهد تا به مرز کارایی برسد. کارایی تکنیکی در این حالت را با TE_i نشان می‌دهیم و این کارایی برای واحد p بصورت زیر محاسبه می‌گردد [۲۱]:

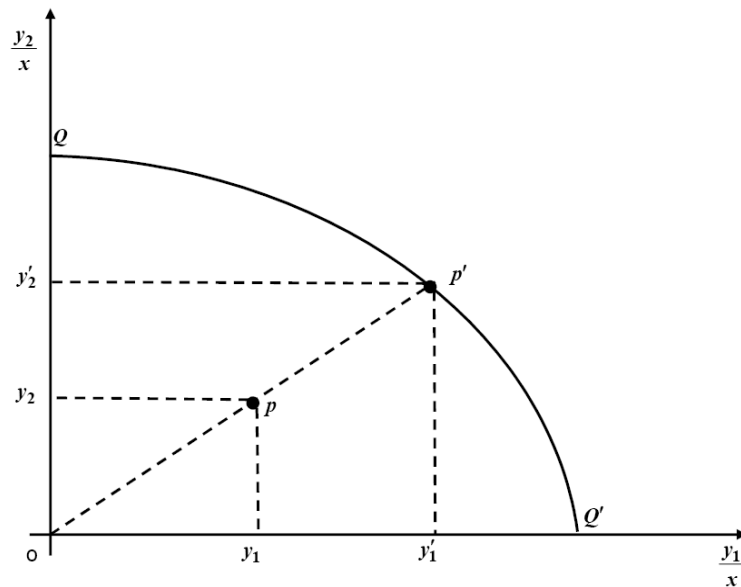
$$TE_i = \frac{op'}{op}$$



شکل ۲-۱: کارایی در حالت ورودی محور

۲-۳-۲ اندازه‌گیری کارایی در حالت خروجی محور

مفهوم خروجی محور این است که به چه میزان باید خروجی‌ها را با ثابت نگه داشتن میزان ورودی‌ها افزایش داد تا واحد مورد نظر به مرز کارایی برسد. برای درک بهتر مفهوم اندازه‌گیری خروجی محور فرض می‌کنیم که واحدهایی در اختیار داریم که دو نوع خروجی y_1 و y_2 را برای ایجاد یک ورودی x استفاده می‌کنند. شکل (۲-۲) این مفهوم را نشان می‌دهد. منحنی QQ' مکان هندسی تمام ترکیبات ممکن از خروجی‌ها است که ورودی یکسان تولید می‌کنند. این منحنی نسبت به مبدأ مقعر می‌باشد که نشان دهنده‌ی محدوده‌ی بالای خروجی‌ها است. یعنی حداکثر سطح ممکن برای افزایش خروجی‌ها را نشان می‌دهد. بنابراین واحدهایی که روی این منحنی قرار می‌گیرند، واحدهای کارا می‌باشند. زیرا با حداقل ممکن مصرف ورودی x به حداکثر خروجی دست یافته‌اند. همان‌طور که بیان شد می‌توان نتیجه گرفت که واحد p غیرکارا است. واحد p برای کارا شدن نیاز دارد که میزان خروجی y_1 خود را به میزان خروجی y_1' و میزان خروجی y_2 خود را به میزان خروجی y_2' افزایش دهد تا به مرز کارایی برسد.



شکل ۲-۲: کارایی در حالت خروجی محور

کارایی تکنیکی در این حالت را با TE_o نشان می‌دهیم و این کارایی برای واحد p بصورت زیر محاسبه می‌گردد [۲۱]:

$$TE_o = \frac{op}{op'}$$

۳-۳-۲ اندازه‌گیری کارایی در تحلیل پوششی داده‌ها

اندازه‌گیری کارایی بر تئوری تولید استوار است. در این تئوری یک شرکت یا یک سازمان یا یک DMU بعنوان یک سیستم تولیدی تلقی می‌شود، که برای ایجاد محصول (خروجی)، منبع (ورودی) را مصرف می‌کند.

$$\text{خروجی} \mapsto DMU \mapsto \text{ورودی}$$

تعریف متداول کارایی، پایه و اساس مدل‌های DEA است. این تعریف بصورت زیر است:

$$k \text{ کارایی واحد } k = \frac{(DMU_k) \text{ خروجی واحد } k}{(DMU_k) \text{ ورودی واحد } k} \quad k = 1, 2, \dots, n$$

اما تعریف فوق تنها برای محاسبه کارایی واحدی مناسب است که دارای یک ورودی و یک خروجی باشد و همچنین در این تعریف نمی‌توان واحدهای مختلف را بطور همزمان دخالت داد. برای رفع این مشکل فارل در سال ۱۹۵۷ یک روش اندازه‌گیری جدید ارائه داد که مبنای تمام مدل‌های DEA قرار گرفت. این روش تمام ورودی‌ها و خروجی‌ها را در برمی‌گیرد و عملکرد هر واحد در مقایسه با یک واحد با بهترین عملکرد تحلیل و ارزیابی می‌شود. روش فارل این را بازگو می‌کند که: یک واحد DMU تا چه میزان می‌تواند خروجی‌هایش را با بهبود کارائیش، افزایش دهد بدون اینکه به منبع جدید نیاز باشد. فارل با در نظر گرفتن وزن‌هایی برای ورودی‌ها و خروجی‌های یک DMU، توانست مدلی با چند ورودی و یک خروجی ارائه دهد. اما هنوز مشکل دخالت چند خروجی در مدل وجود داشت که این مسئله نیز توسط چارنز و همکارانش حل شد. تعریف کارایی که از آن در مدل‌های DEA استفاده می‌شود، بدین صورت است:

$$k \text{ کارایی واحد } k = \frac{\text{مجموع وزنی خروجی‌های واحد } (DMU_k)}{\text{مجموع وزنی ورودی‌های واحد } (DMU_k)} \quad k = 1, 2, \dots, n$$

هر چند مؤلفه‌های بردار ورودی و خروجی یک واحد می‌توانند مقادیر صفر داشته باشند، ولی خود بردار داده نمی‌تواند صفر باشد. یعنی باید حداقل یک ورودی یا خروجی مخالف صفر داشته باشد. این فرض از اینکه یک واحد بدون خروجی داشته باشیم و یا تصور اینکه یک واحد چیزی را به واحدهای دیگر منتقل نمی‌کند، جلوگیری می‌کند. این یک فرض استاندارد در DEA می‌باشد. DMU‌هایی که در مدل‌های DEA مورد استفاده قرار می‌گیرند دارای خصوصیات ورودی و خروجی مشابه می‌باشند و این خصوصیات منعکس کننده فعالیت‌های آن DMU است و این تشابه خصوصیات، اندازه‌گیری با معنی از کارایی نسبی را بدست می‌دهد. حال اگر وزن‌های متناظر با خروجی j —ام را با u_j و وزن‌های متناظر با ورودی i را با v_i نمایش دهیم. در اینصورت بهره‌وری واحد تصمیم‌گیری k بصورت زیر محاسبه

می‌شود:

$$E_k = \frac{\sum_j u_j O_{jk}}{\sum_i v_i I_{ik}} \quad k = 1, 2, \dots, n \quad (1-2)$$

I_k, O_k به ترتیب خروجی و ورودی واحد تصمیم‌گیری k می‌باشند [۵۴].

۴-۳-۲ وزن‌ها و محدودیت وزنی

در قسمت قبل وزن‌های متناظر با خروجی j را با u_j و وزن‌های متناظر با ورودی i را با v_i در نظر گرفتیم و کارایی را بصورت (۱-۲) بیان کردیم. اما مسئله اصلی، تخصیص این وزن‌هاست. بعنوان یک ایده می‌توان از یک مجموعه ثابت برای وزن‌های ورودی‌ها و خروجی‌ها استفاده کرد. اما اینگونه مجموعه‌ها اختیاری است و ممکن است از نقطه نظر برخی تخصیص وزن به شاخص‌ها غیر منصفانه جلوه کند (از نقطه نظر آن‌هایی که در این مجموعه نمره‌ی (وزن) کمتری دریافت می‌کنند) و از طرفی در صورتی که وزن‌ها یکسان باشند ممکن است یک DMU که عملاً غیر کارا است (دارای خروجی j یا ورودی i با ارزش کم) بعنوان یک DMU با ارزش بالاتر برای آن خروجی یا ورودی طبقه‌بندی شود. برای رفع این مشکل چارنز و همکارانش پیشنهاد دادند، که هر DMU اجازه دارد وزنش را خودش انتخاب کند. اما باید قوانینی برای انتخاب وزن‌ها توسط خود DMU وجود داشته باشد. در غیر این صورت باز هم ایده ما غیر منصفانه خواهد بود. اولین قانون این است که بعد از اینکه یک DMU وزن‌های خود را انتخاب کرد، این وزن‌ها برای سایر DMU‌ها استفاده خواهد شد. قانون دیگر این است که به هر حال وزن‌ها دارای یک میزان حداکثر خواهند بود و مقدار آنها محدود خواهد شد. امتیاز کارایی هر DMU نمی‌تواند از یک مقدار مثبت و ثابت (L) تجاوز کند زیرا امتیاز کارایی هر واحد در مقایسه با سایر واحدها محاسبه می‌شود. این

قانون علاوه بر نرمالیزه کردن 1° و جلوگیری از ابتدال، شرایط آسان و شایسته‌ای را برای اینکه یک DMU کارا باشد، فراهم می‌کند. قانون دیگر این است که وزن‌ها نمی‌توانند صفر شوند زیرا در غیر این صورت اگر دو DMU که دارای ارزش ورودی‌ها و خروجی‌های یکسان به جز در یکی از ورودی‌ها (خروجی‌ها) باشند، با انتخاب وزن صفر برای آن ورودی (خروجی) دارای امتیاز کارایی یکسان خواهند شد. برای رفع این مشکل چارنز و همکارانش استفاده از مقدار غیرارشمیدسی ϵ (مقدار ثابت، مثبت و کوچک) را پیشنهاد داده‌اند. حال بر اساس آنچه گفته شد مدل اصلی DEA را ارائه می‌دهیم. اگر ماکزیمم عبارت زیر را بدون هیچگونه محدودیتی بیابیم، در این صورت مسئله نامحدود خواهد شد

$$E_k = \frac{\sum_j u_j O_{jk}}{\sum_i v_i I_{ik}}.$$

به همین خاطر محدودیتی به مسئله اضافه می‌شود تا کارایی حاصل از این وزن‌ها برای تمام واحدها در یک بازه معین $[0, L]$ قرار گیرد. در اغلب موارد، L را برابر یک در نظر می‌گیرند تا کارایی واحدها در بازه $[0, 1]$ قرار گیرد. از طرفی فرض غیر صفر بودن وزن‌ها را با استفاده از مقدار غیرارشمیدسی ϵ لحاظ می‌کنیم. بدلیل مثبت و کوچک بودن ϵ ، در محاسبات امتیاز کارایی نسبی خللی وارد نمی‌کند و از طرفی همانطور که گفته شد تحلیل و طبقه‌بندی DMUها را واقع بینانه‌تر می‌سازد. مدل زیر را خواهیم داشت:

$$\begin{aligned} \max \quad & \frac{\sum_j u_j O_{jk}}{\sum_i v_i I_{ik}} \\ s.t \quad & \frac{\sum_j u_j O_{js}}{\sum_i v_i I_{is}} \leq 1 \quad \forall s \\ & u_j, v_i \geq \epsilon \quad \forall i, j \end{aligned} \quad (2-2)$$

مدل فوق مدل غیرخطی است که تابع هدف آن بصورت کسری از مجموع وزنی خروجی و ورودی است. برای تبدیل مدل‌های فوق به مدل‌های خطی (LP) از روشی که توسط چارنز و کوپر در سال

1° نرمالیزه شده‌ی بردار $A = (a_1, a_2, \dots, a_n)$ به شکل $A_n = (\frac{a_1}{\sum_{i=1}^n |a_i|}, \frac{a_2}{\sum_{i=1}^n |a_i|}, \dots, \frac{a_n}{\sum_{i=1}^n |a_i|})$ می‌باشد.

۱۹۶۳ ارائه شد، استفاده می‌کنیم. اگر قرار دهیم $\frac{1}{t} = \sum_i v_i I_{ik}$ ، چون $t \geq 0$ خواهیم داشت:

$$\begin{aligned} \max \quad & \sum_j u_j t O_{jk} \\ \text{s.t.} \quad & \sum_j t u_j O_{js} - \sum_i t v_i I_{is} \leq 0 \quad s = 1, \dots, K \\ & t u_j \geq \epsilon \quad s = 1, \dots, J \\ & t v_i \geq \epsilon \quad s = 1, \dots, I \end{aligned} \quad (3-2)$$

با قرار دادن $\bar{u}_j = t u_j$ برای هر $j = 1, \dots, J$ و $\bar{v}_i = t v_i$ برای هر $i = 1, \dots, I$ خطی شده مدل

(۲-۲) که به مدل CCR در تحلیل پوششی داده‌ها معروف است به شکل زیر تبدیل می‌شود:

$$\begin{aligned} \max \quad & \sum_j \bar{u}_j O_{jk} \\ \text{s.t.} \quad & \sum_i \bar{v}_i I_{ik} = 1 \\ & \sum_j \bar{u}_j O_{js} - \sum_i \bar{v}_i I_{is} \leq 0 \quad s = 1, 2, \dots, K \\ & \bar{u}_j, \bar{v}_i \geq \epsilon \quad \forall i, j \end{aligned} \quad (4-2)$$

واحد مورد بررسی را کارا گوئیم اگر مقدار تابع هدف برابر یک گردد و در غیر این صورت واحد را ناکارا گوئیم [۵۴، ۸].

از این پس به دلیل سهولت در نمایش از مدل (۴-۲) بدون علامت $(-)$ بر روی وزن‌ها استفاده می‌کنیم.

حال اگر تابع هدف مدل (۴-۲) را با w_k نشان دهیم با تعریف پارامتر $d_k = 1 - w_k$ به عنوان انحراف واحد k از کارایی، داریم [۲۲]:

$$\begin{aligned} \max \quad & w_k = 1 - d_k \\ \text{s.t.} \quad & \sum_i v_i I_{ik} = 1 \quad \forall k \\ & \sum_j u_j O_{jk} + d_k = 1 \quad \forall k \\ & \sum_j u_j O_{js} - \sum_i v_i I_{is} \leq 0 \quad \forall s; s \neq k \\ & u_j, v_i \geq \epsilon \quad \& \quad d_k \geq 0 \quad \forall k \end{aligned} \quad (5-2)$$

حال اگر بخواهیم مجموع کارایی یک مجموعه از واحدها را محاسبه کنیم بدین شکل عمل کرده که کارایی تک تک واحدها را محاسبه و مقدار کارایی آنها را با هم جمع می‌کنیم مدل این مسئله به شکل زیر است:

$$\begin{aligned}
 \max \quad & \sum_k w_k = \sum_k (1 - d_k) \\
 s.t \quad & \sum_i v_{ki} I_{ik} = 1 \quad \forall k \\
 & \sum_j u_{kj} O_{jk} + d_k = 1 \quad \forall k \quad (2-6) \\
 & \sum_j u_{kj} O_{js} - \sum_i v_{ki} I_{is} \leq 0 \quad \forall s \neq k \\
 & u_{kj}, v_{ki} \geq \epsilon \quad \forall k, j, i \\
 & d_k \geq 0 \quad \forall k
 \end{aligned}$$

۵-۳-۲ مثال عددی

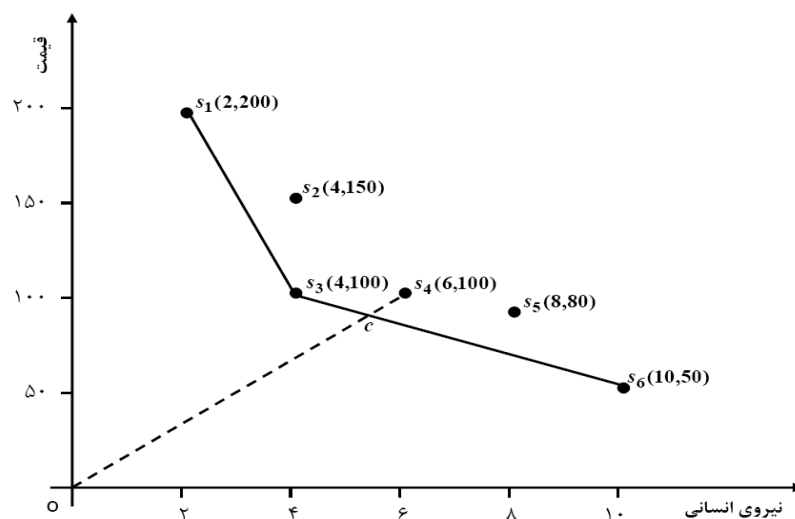
با یک مثال این بخش را به پایان می‌رسانیم [۶۰]. در این مثال می‌خواهیم شش رستوران را از لحاظ کارایی با هم مقایسه کنیم. جدول (۱.۲) دو ورودی مهم شش رستوران شامل نیروی انسانی و قیمت مواد اولیه را برای فروش صد عدد غذای یک نفره نشان می‌دهد. معمولاً خروجی بین واحدهای خدماتی متفاوت است ولی در این مثال تعمداً خروجی را یکسان نموده‌ایم تا بتوانیم مسئله را از طریق ترسیمی نیز حل کرده و نتایج آن را نشان دهیم. همان طور که در شکل (۲-۳) مشاهده می‌شود واحدهای خدماتی s_1, s_3, s_6 در نقاطی قرار گرفته‌اند که اگر به یکدیگر متصل شوند تشکیل یک پوشش محدب می‌دهند و نقاط s_2, s_4, s_5 را در بر می‌گیرند. به همین دلیل است که این روش را روش آنالیز پوشش داده‌ها نامیده‌اند زیرا پوشش محدبی که از وصل کردن نقاط کارآمد بدست می‌آید پакتی را تشکیل می‌دهد که نقاط کارآمد در گوشه‌ها و مرزی آن قرار می‌گیرند و نقاط غیر کارآمد در داخل این پاکت محبوس

جدول ۱.۲: داده‌های ورودی مثال

واحد خدماتی	تعداد غذای فروخته شده	نیروی انسانی (نفر ساعت)	قیمت مواد اولیه (واحد پول)
s_1	۱۰۰	۲	۲۰۰
s_2	۱۰۰	۴	۱۵۰
s_3	۱۰۰	۴	۱۰۰
s_4	۱۰۰	۶	۱۰۰
s_5	۱۰۰	۸	۸۰
s_6	۱۰۰	۱۰	۵۰

می‌شوند.

پس از حل مدل‌های DEA با داده‌های جدول (۱.۲) نتایج جدول (۲.۲) حاصل می‌شود. جدول (۳.۲) محاسبه‌ی مقدار ورودی مازاد مصرفی توسط واحد s_4 را نشان می‌دهد. این جدول شامل محاسبات مربوط به یک واحد فرضی مانند c است که ترکیبی از ورودی‌های واحدهای مرجع s_3, s_6 است. همانطور که در شکل (۳-۲) می‌بینیم واحد فرضی c بر روی تقاطع خط واصل از مبدأ به s_4 و خط واصل s_3, s_6 قرار دارد. از مختصات نقطه‌ی c در می‌یابیم که واحد s_4 در حال حاضر $۰/۷$ نیروی انسانی و $۱۱/۱$ واحد پول اولیه‌ی زیادی مصرف می‌کند. هر واحد غیر کارآمد یک مجموعه واحد کارآمد به عنوان مرجع خواهد داشت که در شکل (۳-۲) واحد غیر کارآمد s_4 در نزدیکی خط واصل مجموعه‌ی کارآمد s_3, s_6 قرار گرفته‌اند ضمناً برداری که از مرکز به s_4 وصل شده خط واصل s_3, s_6 را در نقطه c قطع می‌کند پس واحد مجازی که برای کارا کردن s_4 در نظر می‌گیریم واحد فرضی c است که توسط s_3, s_6 ساخته می‌شود. سهم هر کدام از دو واحد مرجع s_3, s_6 در ساخت c از حل مدل به دست می‌آید که به ترتیب $(۰/۲۲۲۲)$ و $(۰/۷۷۷۸)$ می‌باشد. مقادیر v_1, v_2 که به ترتیب در رابطه با ضرائب ورودی



شکل ۲-۳: مرز کارایی برای مثال عددی

جدول ۲.۲: نتایج حاصل از حل مثال

واحد خدماتی	مجموعه‌ی مرجع	ضریب نیروی انسانی (v_1)	ضریب مواد اولیه (v_2)	کارایی
s_1	—	۰/۱۶۶۷	۰/۰۰۳۳	۱/۰۰۰
s_2	$\{s_1(۰/۲۸۵۷), s_3(۰/۷۱۴۳)\}$	۰/۱۴۲۸	۰/۰۰۲۸	۰/۸۵۷
s_3	—	۰/۰۶۲۵	۰/۰۰۷۵	۱/۰۰۰
s_4	$\{s_3(۰/۷۷۷۸), s_6(۰/۲۲۲۲)\}$	۰/۰۵۵۵	۰/۰۰۶۷	۰/۸۸۹
s_5	$\{s_3(۰/۴۵۴۵), s_6(۰/۵۴۵۴)\}$	۰/۰۵۶۸	۰/۰۰۶۸	۰/۹۰۱
s_6	—	۰/۰۶۲۵	۰/۰۰۷۵	۱/۰۰۰

جدول ۳.۲: مشخصات واحد الگو

اضافه مصرف	s_4	واحد مجازی	مجموعه‌ی مرجع	ورودی و خروجی
		c	s_3 s_6	
۰	۱۰۰	۱۰۰	$(0/7778) \times 100 + (0/2222) \times 100 =$	غذا
۰/۷	۶	۵/۳	$(0/7778) \times 4 + (0/2222) \times 100 =$	کارگر
۱۱/۱	۱۰۰	۸۸/۹	$(0/7778) \times 100 + (0/2222) \times 50 =$	مواد

های نیروی انسانی و مواد اولیه هستند نشان می‌دهند که اگر یک واحد از ورودی مربوطه کم کنیم چه مقدار به کارایی واحد مربوطه اضافه می‌شود برای واحد s_4 اگر یک واحد از نیروی انسانی کم کنیم ($0/5550$) به مقدار کارایی آن افزوده خواهد شد. بنابراین برای اینکه واحد s_4 کارآمد شود باید دو واحد از نیروی انسانی آن کم کنیم زیرا در این صورت $0/111 = 2(0/0555)$ به مقدار کارایی فعلی آن که $0/88$ است اضافه می‌شود و در این صورت کارایی آن به یک می‌رسد.

۴-۲ تحلیل پوششی داده‌های بازه‌ای

مدل‌هایی که در بخش‌های قبل معرفی شدند، واحدهای تصمیم‌گیری را در نظر می‌گیرند که از اعداد قطعی بعنوان ورودی‌ها و خروجی‌هایشان استفاده می‌کنند. ولی در عمل حالت‌های زیادی وجود دارد که نمی‌توان آنها را بصورت قطعی بیان کرد. مثلاً، میزان آلودگی حاصل از آلاینده‌های یک کارخانه، کیفیت کالای تولیدی یک شرکت، هزینه‌های یک طرح و ... را نمی‌توان با اعداد قطعی بیان کرد. لذا ضروری است در این حالات مدل‌های فوق طوری اصلاح شوند که بتوان از اعداد فازی و بازه‌ای نیز در آنها استفاده کرد. فرض کنید در مدل (CCR) ورودی‌ها و خروجی‌ها به صورت قطعی معلوم نباشند و

به شکل بازه‌ای مانند زیر مطرح شوند:

$$\begin{aligned}\tilde{O}_{jk} &= (O_{jk}, \underline{O}_{jk}, \overline{O}_{jk}) & \tilde{I}_{ik} &= (I_{ik}, \underline{I}_{ik}, \overline{I}_{ik}) & \underline{I}_{ik} &\geq 0, \underline{O}_{jk} \geq 0 \\ \tilde{O}_{js} &= (O_{js}, \underline{O}_{js}, \overline{O}_{js}) & \tilde{I}_{is} &= (I_{is}, \underline{I}_{is}, \overline{I}_{is}) & \underline{I}_{is} &\geq 0, \underline{O}_{js} \geq 0\end{aligned}$$

که در آنها مؤلفه‌های دوم و سوم با اندیس k به ترتیب کران پایین و کران بالا برای واحد مورد بررسی و با اندیس s برای تمام واحدها می‌باشند. در این صورت مدل (CCR) به شکل زیر خواهد بود:

$$\begin{aligned}\theta_k &= \max \sum_j u_j [\underline{O}_{jk}, \overline{O}_{jk}] \\ s.t. & \\ & \sum_i v_i [\underline{I}_{ik}, \overline{I}_{ik}] = 1 \\ & \sum_j u_j [\underline{O}_{js}, \overline{O}_{js}] - \sum_i v_i [\underline{I}_{is}, \overline{I}_{is}] \leq 0 \\ & u_j, v_i \geq \epsilon\end{aligned} \quad (7-2)$$

که به یک مدل برنامه ریزی غیر خطی (غیر محدب) با پارمترهای قطعی v_i, u_j و پارمترهای بازه‌ای $I_{is}, I_{ik}, O_{js}, O_{jk}$ تبدیل می‌شود. مقدار بهینه‌ی مدل فوق بصورت یک بازه می‌باشد. یعنی مقدار بهینه‌ی θ_k^* را بصورت $[\theta_k^L, \theta_k^U]$ داریم. با استفاده از روش نورا^{۱۱} [۵۹] کران‌های پایین و بالای بازه برای DMU_k به صورت زیر معرفی می‌شود.

$$\begin{aligned}\theta_k^L &= \max \sum_j u_j \underline{O}_{jk} \\ s.t. & \\ & \sum_i v_i \overline{I}_{ik} = 1 \\ & \sum_j u_j \overline{O}_{js} - \sum_i v_i \underline{I}_{is} \leq 0 \\ & \sum_j u_j \underline{O}_{jk} - \sum_i v_i \overline{I}_{ik} \leq 0 \\ & u_j, v_i \geq 0\end{aligned} \quad (8-2)$$

$$\begin{aligned}
\theta_k^U &= \max \sum_j u_j \bar{O}_{jk} \\
s.t. \quad & \sum_i v_i \underline{I}_{ik} = 1 \\
& \sum_j u_j \underline{O}_{js} - \sum_i v_i \bar{I}_{is} \leq 0 \\
& \sum_j u_j \bar{O}_{jk} - \sum_i v_i \underline{I}_{ik} \leq 0 \\
& u_j, v_i \geq 0
\end{aligned} \tag{۹-۲}$$

مدل (۸-۲) مؤید این نتیجه است که DMU_k در بدترین شرایط قرار دارد، یعنی ورودی‌های آن در کران بالا و خروجی‌های آن در کران پایین می‌باشند و سایر DMUها در ایده‌آل‌ترین حالت قرار دارند.

قضیه ۱.۲ اگر $\theta_k^{U*}, \theta_k^{L*}$ مقادیر بهینه مدل‌های (۸-۲) و (۹-۲) باشند آنگاه $\theta_k^{L*} \leq \theta_k^{U*}$ [۹].

تعریف ۴.۲ DMU_k را کارآمد قوی گویند اگر $\theta_k^{L*} = \theta_k^{U*} = 1$.

DMU_k را کارا گویند اگر $\theta_k^{U*} = 1, \theta_k^{L*} < 1$.

DMU_k را ناکارا گویند اگر $\theta_k^{U*} < 1$.

در ادامه متغیری را در بازه‌ی هر شاخص تعریف می‌کنیم که اجازه داشته باشد کل مقادیر یک بازه را بپذیرد و با وارد کردن این متغیر در مدل این اختیار به پردازشگر داده می‌شود که هر شاخص مربوط به واحد تحت ارزیابی مقداری از یک بازه را بپذیرد که در نهایت آن واحد بیشترین کارایی ممکن را داشته باشد. حال هر متغیر بازه‌ای را با متغیرهای λ_{jk}, t_{ik} که $0 \leq \lambda_{jk}, t_{ik} \leq 1$ به صورت ترکیب خطی از حد بالا و حد پایین آن می‌نویسیم.

$$\underline{O}_{jk} \leq O_{jk} \leq \bar{O}_{jk} \implies O_{jk} = \lambda_{jk} \bar{O}_{jk} + (1 - \lambda_{jk}) \underline{O}_{jk} \implies O_{jk} = \underline{O}_{jk} + \lambda_{jk} (\bar{O}_{jk} - \underline{O}_{jk})$$

به طریقی مشابه داریم

$$O_{js} = \underline{O}_{js} + \lambda_{js} (\bar{O}_{js} - \underline{O}_{js}).$$

برای ورودی‌های داریم

$$\underline{I}_{ik} \leq I_{ik} \leq \bar{I}_{ik} \implies I_{ik} = t_{ik}\bar{I}_{ik} + (1 - t_{ik})\underline{I}_{ik} \implies I_{ik} = \underline{I}_{ik} + t_{ik}(\bar{I}_{ik} - \underline{I}_{ik}).$$

و به طریقی مشابه داریم

$$I_{is} = \underline{I}_{is} + t_{is}(\bar{I}_{is} - \underline{I}_{is}).$$

بعد از جای گذاری روابط بالا در مدل (۷-۲) خواهیم داشت:

$$\begin{aligned} \max \quad & \sum_j u_j (\underline{Q}_{jk} + \lambda_{jk}(\bar{O}_{jk} - \underline{Q}_{jk})) \\ s.t \quad & \\ & \sum_i v_i (\underline{I}_{ik} + t_{ik}(\bar{I}_{ik} - \underline{I}_{ik})) = 1 \\ & \sum_j u_j (\underline{Q}_{js} + \lambda_{js}(\bar{O}_{js} - \underline{Q}_{js})) - \sum_i v_i (\underline{I}_{is} + t_{is}(\bar{I}_{is} - \underline{I}_{is})) \leq 0 \\ & 0 \leq \lambda_{jk}, \lambda_{js}, t_{ik}, t_{is} \leq 1 \\ & u_j, v_i \geq \epsilon \quad \forall i, j \end{aligned} \quad (10-2)$$

که پس از ساده کردن مدل (۱۰-۲) بصورت زیر نوشته می‌شود.

$$\begin{aligned} \max \quad & \sum_j [u_j \underline{Q}_{jk} + u_j \lambda_{jk}(\bar{O}_{jk} - \underline{Q}_{jk})] \\ s.t \quad & \\ & \sum_i [v_i \underline{I}_{ik} + v_i t_{ik}(\bar{I}_{ik} - \underline{I}_{ik})] = 1 \\ & \sum_j [u_j \underline{Q}_{js} + u_j \lambda_{js}(\bar{O}_{js} - \underline{Q}_{js})] - \sum_i [v_i \underline{I}_{is} + v_i t_{is}(\bar{I}_{is} - \underline{I}_{is})] \leq 0 \\ & 0 \leq \lambda_{jk}, \lambda_{js}, t_{ik}, t_{is} \leq 1 \\ & u_j, v_i \geq \epsilon \quad \forall i, j \end{aligned} \quad (11-2)$$

مدل (۱۱-۲) به خاطر تولید شدن متغیرهای $u_j \lambda_{js}$, $u_j \lambda_{jk}$ برای خروجی‌ها و $v_i t_{is}$, $v_i t_{ik}$ برای ورودی

ها همچنان غیرخطی است. برای خطی کردن مدل، متغیرهای q_{ik} , q_{is} , p_{jk} , p_{js} را به شکل زیر در نظر

گرفته و در مدل جای گذاری می‌کنیم:

$$\left\{ \begin{array}{l} u_j \lambda_{js} = p_{js} \\ u_j \lambda_{jk} = p_{jk} \end{array} \right. \implies 0 \leq p_{js} \leq u_j \quad \& \quad \left\{ \begin{array}{l} v_i t_{is} = q_{is} \\ v_i t_{ik} = q_{ik} \end{array} \right. \implies 0 \leq q_{is} \leq v_i$$

بنابراین پس از جای‌گذاری این متغیرها در مدل (۱۱-۲) داریم:

$$\begin{aligned}
 \max \quad & \sum_j [u_j \underline{Q}_{jk} + p_{jk} (\overline{O}_{jk} - \underline{Q}_{jk})] \\
 s.t \quad & \\
 & \sum_i [v_i \underline{I}_{ik} + q_{ik} (\overline{I}_{ik} - \underline{I}_{ik})] = 1 \\
 & \sum_j [u_j \underline{Q}_{js} + p_{js} (\overline{O}_{js} - \underline{Q}_{js})] - \sum_i [v_i \underline{I}_{is} + q_{is} (\overline{I}_{is} - \underline{I}_{is})] \leq 0 \\
 & p_{js} - u_j \leq 0 \quad \forall j, s \\
 & q_{is} - v_i \leq 0 \quad \forall i, s \\
 & p_{js}, q_{is} \geq 0 \quad \forall i, j, s \\
 & u_j, v_i \geq \epsilon \quad \forall i, j
 \end{aligned} \tag{۱۲-۲}$$

با توجه به اینکه اندیس s تمام واحدها را می‌شمارد پس واحد مورد بررسی هم توسط اندیس s شماره می‌شود. مدل (۱۲-۲)، مدل CCR خطی در حالتی است که ورودی‌ها و خروجی‌ها به صورت بازه‌ای مطرح شوند و در آن $\underline{Q}_{js}$ ($\underline{I}_{is}$) حد پایین خروجی (ورودی) $j-i$ ام برای واحد s است و $\overline{O}_{js}$ ($\overline{I}_{is}$) حد بالا خروجی (ورودی) $j-i$ ام برای واحد s است. در مدل (۱۲-۲) اگر حد بالا و پایین هر ورودی و خروجی یکسان در نظر گرفته شود یعنی $\underline{Q}_{js} = \overline{O}_{js}$ و $\underline{I}_{is} = \overline{I}_{is}$ در این صورت ورودی و خروجی‌ها قطعی شده و مدل (۱۲-۲) همان مدل اولیه CCR خواهد شد. این مدل را با در نظر گرفتن مقدار قطعی برای بعضی از ورودی‌ها و خروجی‌ها و مقدار بازه‌ای برای دیگر ورودی و خروجی‌ها می‌توان حل کرد [۴۷، ۹، ۴۹، ۵۶].

۵-۲ تحلیل پوششی داده‌های تصادفی

قبل از وارد شدن به بحث اصلی لازم است بعضی از تعاریف و اصطلاحات بکار رفته در متن اصلی توضیح داده شود. این تعاریف از مرجع [۵۰] گرفته شده است.

متغیر تصادفی: تابعی از فضای نمونه‌ای به مجموعه‌ای اعداد حقیقی است.

تابع توزیع: توزیع احتمالات یک متغیر تصادفی، تابعی است از دامنه آن متغیر بر بازه $[0, 1]$ ، به طوری که احتمال رخ دادن پیشامدهای با مقدار عددی کمتر از آن را نمایش می‌دهد و بصورت دقیق به شکل زیر نمایش داده می‌شود

$$P(X \leq x) = F_X(x).$$

بر اساس اینکه متغیر گسسته یا پیوسته باشد توزیع، گسسته یا پیوسته نام می‌گیرد.

توزیع نرمال: توزیع نرمال ابتدا در سال ۱۷۳۳ در مقاله‌ای توسط آبراهام دموآر^{۱۲} معرفی شد. این مقاله در دومین ویرایش از «قانون اساسی شانسها» در قسمت تخمین قطعی توزیع‌های دوجمله‌ای برای n های بزرگ در سال ۱۷۳۸ مجدداً به چاپ رسید. توزیع نرمال یا توزیع گاوسی یکی از توزیع‌های احتمالی پیوسته است. این توزیع با بردار میانگین و کواریانس آن توصیف می‌شود. به علت شباهت نمودار توزیع نرمال استاندارد به زنگوله به آن انحنای زنگوله‌ای نیز گفته می‌شود. دلیل اهمیت توزیع نرمال از وجود قضیه حد مرکزی ناشی می‌شود. این قضیه می‌گوید هنگامی که تعداد بسیار زیادی متغیر تصادفی با توزیع دلخواه (واریانس محدود) را با هم جمع کنیم و میانگین بگیریم، توزیع نهایی به توزیع نرمال میل می‌کند. به همین خاطر هنگامی که شاهد تاثیر حجمی بسیاری از پدیده‌های تصادفی هستیم، نتیجه نهایی با توزیع نرمال قابل توصیف است. به تجربه ثابت شده است که در دنیای اطراف ما توابع بسیاری از متغیرهای طبیعی از همین تابع پیروی می‌کنند.

توزیع نرمال یک متغیره: متغیر تصادفی X دارای توزیع نرمال است و به آن متغیر تصادفی نرمال اطلاق می‌شود، هرگاه تابع چگالی احتمال آن بصورت زیر باشد:

$$f(x) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2}$$

که در آن μ و σ به ترتیب میانگین و انحراف معیار توزیع نرمالند.

توزیع نرمال استاندارد: توزیع نرمال با $\mu = 0$ و $\sigma^2 = 1$ ، را توزیع نرمال استاندارد می‌نامند.

قضیه ۲.۲ اگر X دارای توزیع نرمال با میانگین μ و انحراف معیار σ باشد، آنگاه $z = \frac{X-\mu}{\sigma}$ توزیع نرمال استاندارد دارد [۵۰].

ماتریس واریانس کواریانس: ماتریسی که روی قطراصلی آن واریانس و بقیه‌ی درایه‌های آن کواریانس

است.

$$\begin{bmatrix} \sigma_1^2 & \sigma_{12} & \cdots & \sigma_{1n} \\ \sigma_{21} & \sigma_2^2 & \cdots & \sigma_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \sigma_{n1} & \sigma_{n2} & \cdots & \sigma_n^2 \end{bmatrix}$$

که در آن $\sigma_i^2 = var(x_i)$ و $\sigma_{ij} = cov(x_i, x_j)$ برای هر $i, j = 1, 2, \dots, n$. ماتریس واریانس -

کواریانس نامیده می‌شود.

یکی از مهمترین نواقص روش تحلیل پوششی داده‌های قطعی متکی بودن بر اطلاعات مربوط به دوره زمانی است که واحدهای تحت بررسی در واقع این دوره زمانی را سپری کرده‌اند. بنابراین نتایجی که حل این مدل به عنوان راهکار به مدیریت ارائه می‌نماید براساس اطلاعات گذشته است. این در حالی است که با در نظر داشتن پویایی عوامل محیطی، تعمیم نتایج مربوط به اطلاعات گذشته جهت تصمیم‌گیری در دوره زمانی آینده نمی‌تواند نتیجه مطلوبی را ایجاد نماید. یکی از راه‌های برطرف‌سازی مشکل فوق ایجاد مدلی است که با در نظر داشتن احتمالات وقوع رویدادها و واردسازی مقادیر پیش‌بینی شده امکان پیش‌بینی کارایی را فراهم سازد. بدین منظور می‌توان از مدلی ریاضی بهره برد که در اصطلاح تحلیل پوششی داده‌های تصادفی^{۱۳} (SDEA) نامیده می‌شود [۴۸]. برخی از مدل‌هایی که با بکارگیری

^{۱۳} Stochastic data envelopment analysis

منطق تحلیل پوششی داده‌ها به منظور پیش‌بینی کارایی قابل استفاده هستند در ادامه بررسی خواهند شد.

۱-۵-۲ مدل تحلیل پوششی داده‌ها محدود شده به قیود تصادفی

این مدل نخستین بار توسط چارلز، کوپر و رودز در سال ۱۹۵۹ با در نظر گرفتن مفاهیمی مانند متغیرهای تصادفی و خطاهای اندازه‌گیری در مدل‌های برنامه ریزی خطی مطرح شد [۴] و بعدها توسط لند^{۱۴}، لاول^{۱۵} و تور^{۱۶} (۱۹۹۳) در قالب مدل (LLT) بسط یافت [۲۵]. این مدل با فرض وجود متغیرهای ورودی و خروجی تصادفی در مدل پوششی تحلیل پوششی داده‌ها مدل نهایی تحلیل پوششی داده‌های تصادفی را ایجاد می‌نماید که دارای محدودیت‌های احتمالی است. هدف اولیه این مدل منظور نمودن خطاهای اندازه‌گیری در مدل پوششی داده‌ها و به دست آوردن مدل تحلیل پوششی داده‌های تصادفی است که این خطاها را در نمرات کارایی واحدها وارد سازد.

۲-۵-۲ مدل رضایت بخشی و مفهوم آن در تحلیل پوششی داده‌ها

پس از مطرح شدن مدل (LLT)، کوپر، هوانگ^{۱۷} و لی^{۱۸} در سال (۱۹۹۶) مدل جدیدی با در نظر داشتن مدل رضایت بخشی سایمون^{۱۹} مطرح نمودند [۷]. این مدل جدید تلفیق مفهوم تصمیم‌گیری رضایت بخشی با مدل‌های تحلیل پوششی داده‌ها محدود شده به قیود تصادفی است که میزان قبول خطاهای تصادفی در محاسبه کارایی واحدهای تصمیم‌گیرنده را به میزان رضایت‌مندی نتایج مدل برای

Land^{۱۴}

Lovell^{۱۵}

Thore^{۱۶}

Huang^{۱۷}

Li^{۱۸}

Simon^{۱۹}

تصمیم‌گیرنده مربوط می‌سازد. این مدل با فرض وجود متغیرهای ورودی و خروجی تصادفی در مدل کسری تحلیل پوششی داده‌ها مدل نهایی تحلیل پوششی داده‌های تصادفی را ایجاد می‌نماید که دارای محدودیت‌های احتمالی است.

۲-۵-۳ مدل بنکر، چارنز و کوپر (BCC) اصلاح شده تصادفی

کوپر، دنگ^{۲۰}، هوانگ و لی در سال (۲۰۰۲) یکی از جدیدترین مدل‌های مطرح شده درباره‌ی تحلیل پوششی داده‌های تصادفی را با تبدیل مدل تحلیل پوششی داده‌ها (بنکر، چارنز و کوپر) به مدل تصادفی نهایی مطرح نمودند [۶]. در این مدل فرض بر این است که با توجه به تصادفی بودن ورودی‌ها و خروجی‌ها، واحدهای تصمیم‌گیرنده نیز نسبت به مقیاس متغیر (افزایشی یا کاهش) بازده خواهند داشت. این مدل نیز اساساً به منظور تعیین کارایی واحدها با فرض وجود بازده نسبت به مقیاس متغیر و با در نظر گرفتن خطاهای اندازه‌گیری تعیین شده است. اگر چه مدل‌های مختلف تحلیل پوششی داده‌های تصادفی جنبه‌های مختلفی از مفهوم تصادفی بودن ورودی‌ها و خروجی‌ها را در بر می‌گیرند ولی دارای شباهت‌های کلی نیز هستند که آنها را از مدل تحلیل پوششی داده‌ها متمایز می‌سازد. جدول (۴.۲) تفاوت‌های بین مدل‌های تحلیل پوششی داده‌ها و تحلیل پوششی داده‌های تصادفی را نشان می‌دهد [۴۸]. همچنان که گفته شد کارایی و ناکارایی یک واحد با توجه به مرز کارایی تعریف می‌گردد لذا چون در مدل‌های SDEA مرز کارایی متأثر از داده‌های تصادفی است بنابراین جهت تشریح کارایی تصادفی ابتدا به بیان نموداری مرز کارایی در مدل‌های SDEA پرداخته و سپس بر اساس این مرز به تعریف کارایی تصادفی خواهیم پرداخت [۱۳]. جهت تشریح بحث شکل (۲-۴) را که بیانگر حالت دو ورودی و یک خروجی می‌باشد در نظر بگیرید:

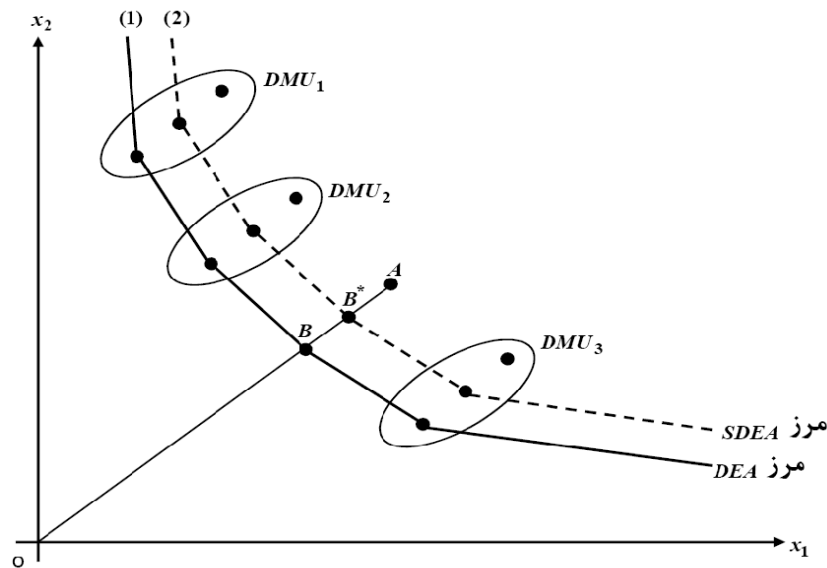
جدول ۴.۲: تفاوت بین مدل های تحلیل پوششی داده‌ها و تحلیل پوششی داده‌های تصادفی

ردیف	مؤلفه‌های مورد مقایسه	ویژگی‌های مدل تحلیل پوششی داده‌ها	ویژگی‌های مدل تحلیل پوششی داده‌های تصادفی
۱	تصادفی بودن	مقادیر ورودی‌ها و خروجی‌ها قطعی هستند.	مقادیر ورودی‌ها و خروجی‌ها تصادفی هستند.
۲	در نظر گرفتن خطای تصادفی	خطای تصادفی در مدل وارد نمی‌شود.	خطای تصادفی در قالب یک جزء تصادفی وارد مدل می‌شود.
۳	پیش بینی کارایی	مدل تنها کارایی مربوط به گذشته را ارائه می‌نماید.	امکان پیش بینی کارایی را فراهم می‌سازد.
۴	مرز کارایی	مرز کارایی نسبت به تغییرات کوچکی در متغیرهای ورودی و خروجی حساسیت زیادی دارا است.	حساسیت کمتری نسبت به تغییرات ایجاد شده در میزان متغیرهای ورودی و خروجی وجود دارد.
۵	تعریف کارایی	سطح خطا یا ریسک مدل ساز (α) صفر منظور می‌گردد.	کارایی با توجه به سطح ریسک مدل ساز (α) تعریف می‌گردد.
۶	همبستگی بین ورودی‌ها و خروجی‌ها	همبستگی بین ورودی‌ها و خروجی‌ها در مدل نهایی منظور نمی‌گردد.	با استفاده از مفهوم واریانس کواریانس همبستگی بین متغیرهای ورودی و خروجی در مدل لحاظ می‌شود.

در DEA قطعی مرز کارایی بوسیله خط پیوسته شماره (۱) نشان داده می‌شود در حالی که در SDEA به علت تصادفی بودن داده‌ها در جستجوی ناحیه اطمینان مربوط به مشاهدات DMU ها می‌باشیم. این فواصل اطمینان بوسیله اشکال بیضوی که پیرامون هر مجموعه از مشاهدات کشیده شده است نشان داده می‌شود. اولسن^{۲۱} و پترسن^{۲۲} [۲۷] مرز کارایی SDEA را با توجه به مرکز نواحی اطمینان به فرم خط شماره (۲) در شکل (۲-۴) نشان دادند طبق شکل مشخص است که برای بعضی مشاهدات کارایی

^{۲۱} Olsen

^{۲۲} Petersen



شکل ۲-۴: مرز کارایی در مدل‌های DEA و SDEA

ممکن است بیش از یک گردد. با توجه به شکل نمره کارایی یک نقطه مانند A در مدل‌های SDEA و DEA بصورت زیر محاسبه می‌گردد:

$$E_{DEA,A} = \frac{oB}{oA} \quad \& \quad E_{SDEA,A} = \frac{oB^*}{oA} \quad (2-13)$$

مقدار کارایی بدست آمده از مدل SDEA یعنی E_{SDEA} معمولاً بیشتر از مقدار کارایی مدل DEA یعنی E_{DEA} می‌باشد. همچنان که از شکل مشخص است بین مرزهای کارایی در مدل‌های SDEA و DEA اختلاف قابل ملاحظه‌ای مشاهده می‌گردد که این فاصله بیانگر میزان خطای تصادفی^{۲۳} با به حساب آوردن تغییرات در عملکرد واحدها است. هر چه واریانس مربوط به داده‌های نمونه بیشتر باشد فاصله اطمینان ارائه شده در مورد داده‌ها بزرگتر می‌شود و براین اساس اختلاف بین نتایج SDEA و DEA بیشتر می‌شود. اکنون به تعریف ریاضی کارایی تصادفی براساس نظر کوپر و همکارانش می‌پردازیم [۴۵].

تعریف ۵.۲. DMU کارای تصادفی است اگر و تنها اگر دو شرط زیر برقرار باشد :

$$(1) \quad \theta^* = 1 \quad \text{که } \theta \text{ بیانگر کارایی است.}$$

(۲) مقدار تمام متغیرهای کمکی باید در تمام جواب‌های بهینه صفر باشد [۴۵].

۲-۵-۴ مدل تحلیل پوششی داده‌های آینده‌نگر

مدل تحلیل پوششی داده‌های آینده‌نگر که به نوعی تلفیق مدل‌های چارنر، کوپر و رودز و مدل رضایت بخشی است، اولین بار در مقاله‌ای توسط سیویوشی^{۲۴} (۱۹۹۹) اشاره گردید [۳۷]. در حالت کسری این مدل که از تبدیل مدل کسری تحلیل پوششی داده‌های مدل (۲-۲) به دست می‌آید به دلیل قبول فرض تصادفی بودن متغیر خروجی، محدودیت‌ها نیز تصادفی می‌گردند. بنابراین، حالت کسری مدل تحلیل پوششی داده‌های آینده‌نگر بدین شرح می‌گردد:

$$\begin{aligned} \max \quad & E(\sum_j u_j \hat{O}_{jk}) \\ \text{s.t.} \quad & \sum_i v_i I_{ik} = 1 \\ & P_r \left[\frac{\sum_j u_j \hat{O}_{js}}{\sum_i v_i I_{is}} \leq \beta_s \right] \geq 1 - \alpha_s \quad \forall s \\ & u_j, v_i \geq 0 \end{aligned} \quad (2-14)$$

در مدل (۲-۱۴)، نمادهای u_j, v_i بیان‌کننده‌ی وزن‌هایی است که به هر یک از ورودی‌ها و خروجی‌ها اختصاص می‌یابد. O_{jk}, I_{ik} به ترتیب بیان‌کننده‌ی i -امین ورودی و j -امین خروجی مربوط به DMU_k است. هم‌چنین نماد P_r بیان‌کننده‌ی احتمال است و از طرفی در صورتی که علامت « $\hat{}$ » در بالای نماد O_{jk} گذاشته شود مشخص می‌گردد که $\hat{O}_{jk}$ یک متغیر تصادفی است. β_s نیز یک مقدار تجویزی از طیف مقادیر ۰ تا ۱ است که بیان‌کننده‌ی سطح کارایی مورد انتظار واحد تصمیم‌گیرنده s -ام است. مقدار α_s به عنوان ریسک‌پذیری تصمیم‌گیرنده در نظر گرفته می‌شود. به عبارت دیگر، $(1 - \alpha_s)$ بیان‌کننده‌ی احتمال رسیدن به سطح مطلوب β_s است. محدودیت دوم در مدل (۲-۱۴) را به صورت زیر بازنویسی

می‌کنیم:

$$P_r \left\{ \sum_j u_j \hat{O}_{js} \leq \beta_s \left(\sum_i v_i I_{is} \right) \right\} \geq 1 - \alpha_s$$

حال در صورتی که از طرفین نامعادله مقدار $\sum_j u_j \bar{O}_{js}$ را کم و بر $\sqrt{V_s}$ تقسیم کنیم رابطه زیر به دست

می‌آید:

$$P_r \left\{ \frac{\sum_j u_j (\hat{O}_{js} - \bar{O}_{js})}{\sqrt{V_s}} \leq \frac{\beta_s \sum_i v_i I_{is} - \sum_j u_j \bar{O}_{js}}{\sqrt{V_s}} \right\} \geq 1 - \alpha_s \quad (15-2)$$

در رابطه (۱۵-۲) $\bar{O}_{js}$ بیان‌کننده‌ی ارزش مورد انتظار $\hat{O}_{js}$ است و V_s بیان‌کننده‌ی ماتریس واریانس-کواریانس است و داراری ماهیت واریانس می‌باشد. در قضیه حد مرکزی چون مخرج کسر، انحراف معیار می‌باشد باید جذر واریانس در مخرج گذاشته شود اما از آنجا که V_s یک ماتریس است در احتمال چندمتغیره از روش‌های مختلفی همچون استفاده از دترمینان ماتریس واریانس-کواریانس و یا ضرب در ماتریس سطری و ستونی ضرایب استفاده می‌گردد.

چون ارزش مورد انتظار که در صورت کسر در قضیه حد مرکزی جایگزین میانگین شده است، ضریب u_j دارد ماتریس واریانس-کواریانس را از چپ در ماتریس سطری u_j ‌ها و از راست در ماتریس ستونی u_j ‌ها ضرب می‌کنیم تا مقداری عددی حاصل گردد. حال متغیر $\hat{Z}_s$ را به صورت زیر تعریف می‌کنیم:

$$\hat{Z}_s = \frac{\sum_j u_j (\hat{O}_{js} - \bar{O}_{js})}{\sqrt{V_s}} \quad s = 1, 2, \dots, n$$

در صورت قبول فرض نرمال بودن برای توزیع احتمال متغیر تصادفی $\hat{O}_{js}$ می‌توانیم چنین استنباط کنیم که متغیر $\hat{Z}_s$ از توزیع نرمال با میانگین یک و واریانس صفر تبعیت می‌کند. با جای‌گذاری رابطه‌ی اخیر در (۱۵-۲) داریم:

$$P_r \left\{ \hat{Z}_s \leq \frac{\beta_s \sum_i v_i I_{is} - \sum_j u_j \bar{O}_{js}}{\sqrt{V_s}} \right\} \geq 1 - \alpha_s \quad (16-2)$$

معکوس شده رابطه (۲-۱۶) به صورت زیر در می‌آید:

$$\frac{\beta_s \sum_i v_i I_{is} - \sum_j u_j \bar{O}_{js}}{\sqrt{V_s}} \geq F^{-1}(1 - \alpha_s) \quad s = 1, 2, \dots, n \quad (2-17)$$

در این جا وجود F بیان کننده‌ی بکارگیری تابع توزیع تجمعی براساس توزیع نرمال و F^{-1} نشان دهنده‌ی تابع معکوس آن است. شکل کلی مدل تحلیل پوششی داده‌های آینده‌نگر با جایگزین کردن رابطه‌ی (۲-۱۷) به صورت زیر در می‌آید:

$$\begin{aligned} \max \quad & E(\sum_j u_j \hat{O}_{jk}) \\ \text{s.t} \quad & \\ & \sum_i v_i I_{ik} = 1 \quad (2-18) \\ & \beta_s \sum_i v_i I_{is} - \sum_j u_j \bar{O}_{js} \geq \sqrt{V_s} F^{-1}(1 - \alpha_s) \quad s = 1, 2, \dots, n \\ & u_j, v_i \geq 0 \end{aligned}$$

حالا می‌توان به منظور تبدیل تابع هدف چنین فرض کرد که مقدار یک متغیر تصادفی $\hat{O}_{jk}$ که در مورد تمام خروجی‌ها صدق می‌کند از طریق رابطه‌ی زیر به دست می‌آید:

$$\hat{O}_{jk} = \bar{O}_{jk} \pm \xi b_{jk} \quad k = 1, 2, \dots, n \quad (2-19)$$

در این جا $\bar{O}_{jk}$ عبارت از مقدار منتظره از خروجی j -ام (میانگین) و b_{jk} عبارت از انحراف معیار آن می‌باشد. ξ بیان کننده‌ی یک متغیر تصادفی است که از توزیع نرمال با میانگین صفر و واریانس σ^2 پیروی می‌کند. در ادامه تشریح می‌کنیم که چگونه می‌توان میانگین مقدار مورد انتظار از خروجی j -ام ($\bar{O}_{jk}$) و انحراف استاندارد آن (b_{jk}) را محاسبه کرد. تحت این مفروضات مقدار V_s چنین می‌شود:

$$\begin{aligned} V_s &= (u_1, u_2, \dots, u_J) \times \begin{bmatrix} b_{1s}^2 \sigma^2 & b_{1s} b_{2s} \sigma^2 & \dots & b_{1s} b_{js} \sigma^2 \\ b_{2s} b_{1s} \sigma^2 & b_{2s}^2 \sigma^2 & \dots & b_{2s} b_{js} \sigma^2 \\ \vdots & \vdots & \ddots & \vdots \\ b_{js} b_{1s} \sigma^2 & \dots & \dots & b_{js}^2 \sigma^2 \end{bmatrix} \times \begin{bmatrix} u_1 \\ u_2 \\ \vdots \\ u_J \end{bmatrix} \quad (2-20) \\ &= \left(\sum_j u_j b_{js} \sigma \right)^2 \end{aligned}$$

پس از اعمال رابطه‌ی (۲۰-۲) در مدل (۱۸-۲) داریم:

$$\begin{aligned} \max \quad & E(\sum_j u_j \hat{O}_{jk}) \\ \text{s.t.} \quad & \sum_i v_i I_{ik} = 1 \\ & \beta_s \sum_i v_i I_{is} - \sum_j u_j \bar{O}_{js} \geq \left(\sum_j u_j b_{js} \sigma \right) F^{-1}(1 - \alpha_s) \quad s = 1, 2, \dots, n \\ & u_j, v_i \geq 0 \end{aligned} \quad (21-2)$$

هم چنان که قبلاً اشاره شد داریم $\hat{O}_{jk} = \bar{O}_{jk} \pm \xi b_{jk}$ پس می‌توان تابع هدف رابطه (۲۱-۲) را به صورت زیر دوباره نویسی کرد:

$$E(\sum_j u_j \hat{O}_{jk}) = E(\sum_j u_j (\bar{O}_{jk} \pm \xi b_{jk})) = E(\sum_j u_j \bar{O}_{jk} \pm \sum_j u_j \xi b_{jk}) = \sum_j u_j \bar{O}_{jk} \quad (22-2)$$

در رابطه (۲۲-۲) به علت اینکه ξ متغیر تصادفی با میانگین صفر و واریانس (δ^2) است در نتیجه داریم:

$$E(\sum_j u_j \xi b_{jk}) = 0$$

تحت این مفروضات یعنی با در نظر گرفتن میانگین صفر و انحراف معیار برابر یک برای ξ مدل تحلیل پوششی داده‌های آینده‌نگر به صورت زیر نوشته می‌شود:

$$\begin{aligned} \max \quad & \sum_j u_j \bar{O}_{jk} \\ \text{s.t.} \quad & \sum_i v_i I_{ik} = 1 \\ & \beta_s \sum_i v_i I_{is} - \sum_j u_j \bar{O}_{js} \geq \left(\sum_j u_j b_{js} \right) F^{-1}(1 - \alpha_s) \quad s = 1, 2, \dots, n \\ & u_j, v_i \geq 0 \end{aligned} \quad (23-2)$$

در اینجا سه نوع برآورد را به یک ارزش مورد انتظار و نیز واریانس مورد انتظار خروجی تبدیل می‌کنیم.

ارزش مورد انتظار خروجی‌ها در این روش به شرح زیر است

$$\bar{O}_{js} = \frac{(OP_{js} + 4ML_{js} + PE_{js})}{6}.$$

واریانس درونی $\hat{O}_{js}$ نیز بصورت زیر است:

$$b_{js}^2 = \frac{(OP_{js} - PE_{js})}{6}$$

که در آنها

ML_{js} : بیان کننده‌ی محتمل‌ترین تخمین از مقدار متغیر خروجی است.

OP_{js} : بیان کننده‌ی تخمین خوش‌بینانه از مقدار متغیر خروجی است.

PE_{js} : بیان کننده‌ی تخمین بدبینانه از مقدار متغیر خروجی است.

مزیت به کارگیری روش تحلیل پوششی داده‌های تصادفی نسبت به روش تحلیل پوششی داده‌ها این است که مدیریت می‌تواند با قبول ریسک متفاوت، کارایی واحدهای تحت سرپرستی خود را پیش‌بینی نماید و به این ترتیب، با در نظر داشتن وضعیت آینده واحدها جهت بهبود عملکرد واحدهای ناکارا اقدام کند. این در حالی است که با در نظر داشتن ماهیت تصادفی خروجی‌های واحدها و به علت اتکای روش تحلیل پوششی داده‌ها بر اطلاعات گذشته، مدیریت نمی‌تواند با تعمیم نتایج به دست آمده از روش‌های قبلی برای آینده واحدهای ناکارا برنامه‌ریزی نماید.

با توجه به کارایی پیش‌بینی شده، مدیریت می‌تواند از طریق شناسایی واحدهای ناکارا، شناسایی واحدهای مرجع این واحدهای ناکارا و هدف‌گذاری برای واحدهای ناکارا در جهت افزایش کارایی این واحدها، قبل از عملکرد واقعی آن‌ها برنامه‌ریزی نماید [۴۸، ۵، ۴۵].

۶-۲ محدودیت‌های وزنی در روش تحلیل پوششی داده‌ها

روش تحلیل پوششی داده‌ها دارای محدودیت‌هایی نیز می‌باشد که از آن جمله می‌توان به عدم کنترل وزن شاخص‌ها و همچنین عدم رضایت مدیریت از اوزان بدست آمده‌ی ورودی‌ها و خروجی‌ها اشاره

کرد. جهت رفع محدودیت‌های ذکر شده، در این جا رویه‌ای آورده شده است که در آن ابتدا نرخ اهمیت ورودی‌ها و خروجی‌ها طبق نظر کارشناسان و با استفاده از تحلیل سلسله مراتبی AHP^{۲۵} بدست می‌آید. این نرخ‌ها به عنوان «وزن نسبی» شاخص‌ها فرض شده و با درویکرد متفاوت به صورت «محدودیت‌های کنترل وزن نسبی» هر شاخص به مدل تحلیل پوششی داده‌ها افزوده می‌شوند. در رویکرد اول، وزن ورودی‌ها و خروجی‌های واحد تحت بررسی به گونه‌ای بدست می‌آیند که وزن نرمالیزه شده آنها با اهمیت نسبی مد نظر کارشناسان به طور دقیق برابر شود. اما در رویکرد دوم وزن شاخص‌ها از مدل به گونه‌ای بدست می‌آیند که «وزن نسبی» آنها با تقریب‌های مختلف با اهمیت مد نظر کارشناسان مطابقت داشته باشد. روش تحلیل پوششی داده‌ها تکنیکی مبتنی بر رویکرد برنامه‌ریزی خطی است و برای هر واحد سازمانی به صورت جداگانه اجرا می‌گردد. این امر در برخی مواقع سبب توزیع غیر واقعی وزن به ورودی‌ها و خروجی‌های مدل می‌گردد. به عبارت دیگر در مدل‌های پایه‌ای تحلیل پوششی داده‌ها دامنه‌ی تغییر متغیرهای وزنی روی یک مجموعه‌ی بیکران نامنفی مجاز شمرده می‌شود تا بدین وسیله اوزان ورودی‌ها و خروجی‌ها به گونه‌ای یافت شوند که واحد مورد ارزیابی نسبت به سایر واحدهای موجود در مجموعه، در بهترین مکان ممکن قرار گیرد. بنابراین ممکن است وزن‌های بدست آمده برای ورودی‌ها و خروجی‌های مشابه از یک واحد به واحد دیگر تفاوت زیادی داشته باشد. این عدم کنترل وزن می‌تواند آسیب‌هایی نظیر موارد زیر را ایجاد کند [۲]:

الف) اوزانی که توسط مدل DEA بدست می‌آید ممکن است همان اوزان مورد نظر مدیریت نباشد (اوزانی که مدیریت با توجه به اهمیت نسبی متغیرها برای آنها در نظر می‌گیرد).

^{۲۵} Analytic hierarchy process

ب) ممکن است یک ورودی یا خروجی مشخص وزن نامناسبی کسب کند به عبارت دیگر ممکن است مدل به خروجی خاص از واحد مورد ارزیابی که مقدار تولیدی (ورودی) آن زیاد است و یا به ورودی خاصی که مقدار مصرفی (خروجی) آن کم است، وزن زیادی تخصیص دهد و از طرف دیگر به خروجی‌هایی با مقدار تولیدی (ورودی) کم و یا به ورودی‌هایی با مقدار مصرفی (خروجی) زیاد، وزن کمتری را اختصاص دهد.

در سال ۱۹۸۸، دایسون^{۲۶} و تاناسولیس^{۲۷} مقاله‌ای را ارائه دادند که در آن در حالت یک ورودی و چند خروجی، وزن‌ها از پایین کراندار شده بودند [۱۱]. رل^{۲۸}، کک^{۲۹} و گلونی^{۳۰} بیان می‌کنند که عدم توجه به موضوع کنترل وزن، عملاً مجاز شمردن تسلط عوامل کم اهمیت در ارزیابی می‌باشد به طوری که واحدهای توانمند در عوامل با اهمیت تر، بعد از واحدهای توانمند در عوامل کم اهمیت تر قرار گرفته و بنابراین نتایج ارزیابی مدل مذکور فاقد اعتبار می‌باشد. همچنین آنها حذف برخی از ورودی‌ها و خروجی‌های ارزیابی عملکرد را که با دقت زیادی هم انتخاب شده‌اند به دلیل تخصیص وزن صفر توسط مدل، عاملی برای بی‌اعتبار بودن مدل می‌دانند. در سال ۱۹۹۱، رل و همکارانش بر لزوم کراندار کردن وزن‌ها در تحلیل پوششی داده‌ها اشاره کردند و روش‌های تجربی برای کراندار کردن وزن‌ها ارائه کردند [۳۲]. در ادامه این کار، در سال ۱۹۹۳ رل و گلونی علاوه بر تکمیل روش‌های ارائه شده قبلی، روش‌های عملی دیگری را نیز مطرح ساختند [۳۳].

محققان در این زمینه سعی کرده‌اند تحلیل پوششی داده‌ها به گونه‌ای بسط داده شود که ضمن سنجش

^{۲۶} Dyson

^{۲۷} Thanassoulis

^{۲۸} Roll

^{۲۹} Cook

^{۳۰} Golany

اندازه کارایی واحدهای سازمانی، محدودیت‌های فوق را نداشته باشد. برای این منظور نظر کارشناسان در تعیین اهمیت هر ورودی در مقایسه با سایر ورودی‌ها و همچنین اهمیت هر خروجی در مقایسه با سایر خروجی‌ها، به صورت اعداد فازی مثلثی با پارامتر α به مدل تحلیل پوششی داده‌ها افزوده شده و مدل مفروض به یک مدل برنامه‌ریزی پارامتری تبدیل می‌شود؛ هر چه مقدار α در این مدل بیشتر باشد «وزن نسبی» بدست آمده از مدل برای ورودی‌ها و خروجی‌ها با اهمیت مدنظر کارشناسان تطابق بیشتری می‌یابد. به منظور کنترل وزن شاخص‌ها، مدل تحلیل پوششی داده‌ها با دو رویکرد زیر بسط داده شده است. در رویکرد اول وزن شاخص‌ها از مدل به گونه‌ای بدست می‌آیند که «وزن نسبی» آنها به طور دقیق با اهمیت مدنظر کارشناسان برابر شود. اما در رویکرد دوم وزن شاخص‌ها از مدل به گونه‌ای بدست می‌آیند که «وزن نسبی» آنها با تقریب‌های مختلف با اهمیت مدنظر کارشناسان مطابقت داشته باشد.

۲-۶-۱ اعمال محدودیت‌های کنترل «وزن نسبی» شاخص‌ها در مدل DEA

در این رویکرد ابتدا وزن مورد نظر مدیریت در رابطه با ورودی‌ها و خروجی‌ها بدست می‌آید. اما این اوزان به طور مستقیم به مدل افزوده نمی‌شود، بلکه وزن بدست آمده برای هر ورودی یا خروجی در مقایسه با سایر ورودی‌ها و خروجی‌ها تلقی شده و سپس این اوزان نسبی به مدل تحلیل پوششی داده‌ها افزوده می‌شوند. به عبارت دیگر در این مدل برخلاف مدل‌های مشابه، اوزان بدست آمده از نظر کارشناسان به عنوان «وزن نسبی» شاخص‌ها (و نه به عنوان وزن قطعی شاخص‌ها) به مدل‌های تحلیل پوششی داده‌ها افزوده می‌شود. مراحل انجام این روش به صورت زیر است:

(۱) نظر مدیریتی کارشناسان در رابطه با وزن ورودی i - $\bar{w}_i$ و همچنین وزن خروجی j - $\bar{w}_j$

با استفاده از تکنیک‌هایی همچون روش تحلیل سلسله مراتبی (AHP) گروهی بدست می‌آیند [۴۴].

این اوزان به عنوان «وزن نسبی» هر ورودی یا خروجی تلقی می‌شوند.

۲) با استفاده از «اوزان نسبی» فوق، برای هر یک از ورودی‌ها و خروجی‌ها، یک «محدودیت کنترل

وزن نسبی» به صورت زیر افزوده می‌گردد:

$$\bar{w}_i = \frac{v_i}{\sum_i v_i} \quad \& \quad \bar{w}_j = \frac{u_j}{\sum_j u_j}$$

که در آنها u_j و v_i وزن خروجی j ام و ورودی i ام می‌باشد. با افزودن این محدودیت‌ها به مدل تحلیل پوششی داده‌ها و حل مدل بدست آمده، وزن ورودی‌ها و خروجی‌ها برای واحد تحت بررسی به گونه‌ای بدست می‌آید که نسبت وزن هر ورودی بر مجموع اوزان ورودی‌ها (وزن نسبی هر ورودی) و نسبت وزن هر خروجی بر مجموع اوزان خروجی‌ها (وزن نسبی هر خروجی)، با اوزان بدست آمده‌ی ورودی‌ها و خروجی‌ها طبق نظر کارشناسان، برابر باشد. به عبارت دیگر این محدودیت‌ها به عنوان «محدودیت‌های کنترل وزن نسبی» به مدل تحلیل پوششی داده‌ها افزوده می‌شوند. جهت افزودن این محدودیت‌ها به مدل‌های تحلیل پوششی داده‌ها می‌بایست ابتدا آنها را به محدودیت‌های خطی تبدیل کرد. به عنوان نمونه در مدل CCR نظر کارشناسان در رابطه با اهمیت نسبی ورودی‌ها و خروجی‌ها به صورت زیر اعمال می‌شود:

$$\begin{aligned} \max \quad & \sum_j u_j O_{jk} \\ \text{s.t} \quad & \sum_i v_i I_{ik} = 1 \\ & \sum_j u_j O_{js} - \sum_i v_i I_{is} \leq 0 \quad \forall s \\ & \bar{w}_i (\sum_i v_i) - v_i = 0 \quad \forall i \\ & \bar{w}_j (\sum_j u_j) - u_j = 0 \quad \forall j \\ & u_j, v_i \geq \epsilon \quad \forall i, j \end{aligned} \quad (24-2)$$

در این مدل $\bar{w}_i, \bar{w}_j$ اعداد حقیقی بوده و به ترتیب وزن بدست آمده‌ی خروجی j -ام و ورودی i -ام بر طبق نظر کارشناسان می‌باشند.

۲-۶-۲ اعمال محدودیت‌های کنترل «وزن نسبی» با تقریب‌های مختلف

در مدل‌های اصلی DEA هیچ وزن از پیش تعیین شده‌ای برای ورودی‌ها و خروجی‌ها وجود ندارد و دامنه‌ی تغییر متغیرهای وزنی روی یک مجموعه‌ی بیکران نامنفی مجاز شمرده می‌شود. بنابراین «وزن نسبی» هر شاخص می‌تواند در بازه‌ی $[0, 1]$ بدست آید. اما در مدل ارائه شده در رویکرد اول، وزن شاخص‌ها می‌بایست به گونه‌ای بدست آیند تا «وزن نسبی» آنها به طور دقیق با اهمیت مدنظر کارشناسان برابر شود. بین دو رویکرد «اوزان نسبی آزاد» و «اوزان نسبی دقیق» می‌توان رویکرد میانه‌ای را یافت به گونه‌ای که «اوزان نسبی» شاخص‌ها با تقریب‌های مختلف با اهمیت مدنظر کارشناسان مطابقت داشته باشد. برای این منظور اوزان بدست آمده از نظر کارشناسان برای ورودی i -ام و خروجی j -ام $(\bar{w}_i, \bar{w}_j)$ به صورت اعداد فازی مثلثی $\tilde{w}_i = (0, \bar{w}_i, 1)$ و $\tilde{w}_j = (0, \bar{w}_j, 1)$ در نظر گرفته می‌شود. با استفاده از روش تبدیل اعداد فازی مثلثی به بازه‌های فازی مبتنی بر α -برش این «اوزان نسبی فازی» به بازه‌های زیر تبدیل می‌شوند:

$$\tilde{w}_i = (0, \bar{w}_i, 1) \Rightarrow [\bar{w}_i \alpha, 1 - (1 - \bar{w}_i) \alpha]$$

$$\tilde{w}_j = (0, \bar{w}_j, 1) \Rightarrow [\bar{w}_j \alpha, 1 - (1 - \bar{w}_j) \alpha]$$

α نشان دهنده‌ی درجه تطابق وزن نسبی بدست آمده از مدل با اوزان بدست آمده از نظر کارشناسان می‌باشد. بنابراین هر چه α بزرگتر باشد نظر کارشناسان در رابطه با اهمیت نسبی شاخص‌ها دقیق‌تر اعمال می‌گردد. از آنجا که «وزن نسبی» هر شاخص در بازه‌ی متناظر با آن شاخص بدست می‌آید.

بنابراین می‌توان حدود «وزن نسبی» ورودی i -ام و خروجی j -ام را به صورت زیر در نظر گرفت:

$$w_i \alpha \leq \frac{v_i}{\sum_i v_i} \leq 1 - (1 - w_i) \alpha$$

$$w_j \alpha \leq \frac{u_j}{\sum_j u_j} \leq 1 - (1 - w_j) \alpha$$

جهت افزودن حدود «وزن نسبی» شاخص‌ها به مدل تحلیل پوششی داده‌ها می‌بایست ابتدا آنها را به

فرم محدودیت‌های خطی نوشت. بنابراین محدودیت مربوط به ورودی i -ام به صورت زیر می‌باشد

$$w_i \alpha \sum_i v_i - v_i \leq 0$$

$$v_i - (1 - \alpha + w_i \alpha) \sum_i v_i \leq 0$$

و محدودیت مربوط به خروجی j -ام نیز به صورت زیر می‌باشد

$$w_j \alpha \sum_j u_j - u_j \leq 0$$

$$u_j - (1 - \alpha + w_j \alpha) \sum_j u_j \leq 0$$

و به مدل تحلیل پوششی داده‌ها افزوده می‌شوند. به عنوان نمونه در مدل (CCR) نظر کارشناسان در

رابطه به اهمیت نسبی ورودی‌ها و خروجی‌ها و با تقریب α ، به صورت زیر اعمال می‌شود:

$$\begin{aligned} \max \quad & \sum_j u_j O_{jk} \\ \text{s.t.} \quad & \sum_i v_i I_{ik} = 1 \\ & \sum_j u_j O_{js} - \sum_i v_i I_{is} \leq 0 \quad \forall s \\ & w_i \alpha (\sum_i v_i) - v_i \leq 0 \quad \forall i \quad (25-2) \\ & w_j \alpha (\sum_j u_j) - u_j \leq 0 \quad \forall j \\ & v_i - (1 - \alpha + w_i \alpha) \sum_i v_i \leq 0 \quad \forall i \\ & u_j - (1 - \alpha + w_j \alpha) \sum_j u_j \leq 0 \quad \forall j \\ & u_j, v_i \geq \epsilon \quad \forall i, j \end{aligned}$$

مدل فوق یک مدل برنامه‌ریزی پارامتری و به فرم مدل‌های (DEA) فازی می‌باشد. با این تفاوت که

مقدار شاخص‌ها معین و قطعی هستند و «وزن نسبی» آنها به صورت نامعین و فازی به مدل افزوده

شده است. بنابراین جهت حل این مدل می‌توان از رویکرد مبتنی بر سطح α استفاده کرد به گونه‌ای که اندازه کارایی هر واحد به ازاء سطوح مختلف α بدست آید. ملاحظه می‌شود که این مدل به ازاء $\alpha = 0$ به مدل کلاسیک (CCR) تبدیل می‌شود. همچنین مدل مذکور به ازاء $\alpha = 1$ به مدل رویکرد اول تبدیل شده و نظر کارشناسان در رابطه با اهمیت نسبی ورودی‌ها و خروجی‌ها به طور دقیق اعمال می‌شود. به طور کلی اعمال نظر کارشناسان در مدل‌های تحلیل پوششی داده‌ها دارای مزایایی می‌باشد که از آن جمله می‌توان به این موارد اشاره کرد، اعمال نظر کارشناسان در مدل به صورت محدودیت‌های کنترل وزن نسبی ورودی‌ها و خروجی‌ها، محاسبه اندازه کارایی واحدها به ازاء مقادیر مختلف α ، افزایش قدرت تفکیک مدل و کاهش تعداد واحدهای کارا و یافتن یک مجموعه‌ی مشترک «وزن نسبی» برای تمام واحدها [۴۳].

فصل ۳

کاربرد تحلیل پوششی داده‌ها در تعیین وزن

رئوس در مسئله‌ی مکانیابی

۳-۱ مقدمه

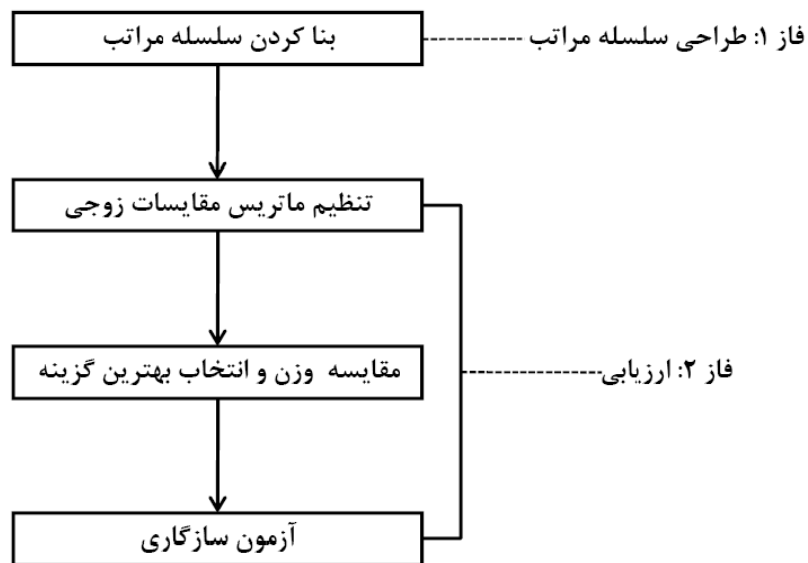
در این فصل بدنال آن هستیم که الگوریتمی را برای تخصیص وزن به رؤس در مسئله‌ی مکانیابی معرفی کنیم. در این الگوریتم از تلفیق روش تحلیل پوششی داده‌ها و فرآیند تحلیل سلسله مراتبی استفاده شده است. البته در اینجا مدل خاصی از تحلیل پوششی داده‌ها که برای تلفیق با AHP مناسب باشد معرفی می‌شود. ایده ترکیب دو روش AHP و DEA جدید نیست. ونگ و همکارانش [۳۸] روش‌های AHP و DEA را برای مسئله‌ی انتخاب سایت ترکیب نمودند. یانگ^۱ و کو^۲ [۴۱] یک روش سلسله مراتبی AHP و DEA را برای سهولت طراحی مسئله مطرح نمودند. حسین زاده و همکارانش [۲۰] یک روش برای تعیین کارایی نسبی واحدهای تصمیم‌گیری با استفاده از ترکیب AHP و DEA پیشنهاد کردند. رمانتان^۳ [۳۰] روشی بنام DEAHP برای تعیین وزن نسبی و نهایی گزینه‌ها پیشنهاد کرد. ونگ و همکارانش [۳۹] یک روش ترکیبی AHP و DEA با ناحیه‌ی اطمینان برای تعیین وزن نسبی و نهایی گزینه‌ها پیشنهاد کردند که محدودیت‌های روش DEAHP را ندارد. مادر این بخش به معرفی روش ترکیبی AHP و DEA که توسط استرن^۴ و همکارانش [۳۶] ارائه شده می‌پردازیم. که در آن از یک رویکرد تلفیقی کمی و کیفی برای وزندهی رؤس استفاده شده است. ماتریس مقایسات زوجی موجود در روش تحلیل سلسله مراتبی بر مبنای حل مدل تحلیل پوششی داده‌ها برای هر جفت از رؤس پُر می‌گردد که این روش، یک روش کمی است اما وزندهی نهایی با استفاده از یک رویکرد کیفی صورت می‌گیرد.

Yang^۱Kuo^۲Ramanathan^۳Stern^۴

۲-۳ فرآیند تحلیل سلسله مراتبی

فرآیند تحلیل سلسله مراتبی در ابتدا در سال ۱۹۷۹ توسط توماس ال ساعتی^۵ [۳۴] معرفی شد اولین کار تخصصی در این زمینه در کتاب ساعتی با عنوان «فرآیند تحلیل سلسله مراتبی» که در سال ۱۹۸۰ منتشر شده، یافت می‌شود. فرآیند تحلیل سلسله مراتبی، بویژه برای تصمیم‌گیری با معیارهای چندگانه مناسب است. این روش، تکنیکی است که برای رتبه‌بندی مجموعه‌ای از گزینه‌ها یا برای انتخاب بهترین، از یک مجموعه گزینه بکار می‌رود. این روش هنگامی که عمل تصمیم‌گیری با چند گزینه رقیب و چند معیار تصمیم‌گیری روبرو است، می‌تواند استفاده گردد. معیارهای مطرح شده می‌توانند کمی و کیفی باشند. اساس تکنیک مذکور، تصمیم‌گیری بر مبنای مقیاسات زوجی است. کاربرد عملی فرآیند تحلیل سلسله مراتبی شامل چهار مرحله اساسی است. چنانچه این مراحل در دوفاز کلی طراحی سلسله مراتب و ارزیابی طبقه‌بندی شود، مرحله اول در فاز طراحی و مرحله بعدی در فاز دوم، یعنی فاز ارزیابی قرار می‌گیرد (شکل ۱-۳). اولین قدم در فرآیند سلسله مراتبی، ایجاد یک نمایش گرافیکی از مسئله است. که در آن هدف، معیارها و گزینه‌ها نشان داده می‌شوند. در رأس سلسله مراتب، هدف کلان و کلی موضوع تصمیم‌گیری و در مراتب پایین‌تر، صفات و معیارهایی قرار گرفته‌اند که به نحوی هریک از آنها در کیفیت هدف تأثیر دارند (چنانچه لازم باشد می‌توان معیارها را به زیرگروه‌های جزئی‌تر تقسیم نمود) و در آخرین سطح گزینه‌ها و انتخاب‌های تصمیم قرار می‌گیرند. نخستین گام در تعیین اولویت‌های عناصر در تصمیم‌گیری، مقایسه‌های دودویی آنهاست، یعنی مقایسه عناصر به صورت جفت جفت با توجه به معیارهای معین، شکل ترجیحی برای انجام مقایسه‌های دودویی، ماتریس می‌باشد. برای فهم بهتر موضوع در شکل (۲-۳) یک سلسله مراتب به نمایش گذاشته شده است و ماتریس مقایسات زوجی

Thomas L Saaty^۵



شکل ۳-۱: مراحل اجرایی AHP

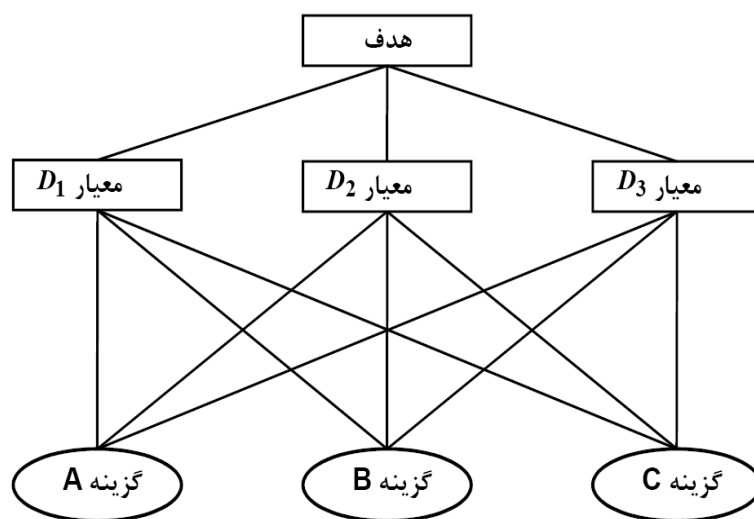
جدول ۱.۳: ماتریس مقایسات زوجی برای تعیین وزن گزینه‌ها برای هر $i = 1, 2, 3$

معیار D_i	A	B	C
A	۱		
B		۱	
C			۱

برای این سلسله مراتب برای تعیین وزن گزینه‌ها در جدول (۱.۳) آمده است. در فرم اصلی فرایند تحلیل سلسله مراتبی جاهای خالی در جدول (۱.۳) توسط کارشناس پرمی‌گردد. برای تعیین وزن هر معیار در تصمیم‌گیری، تصمیم‌گیرنده باید ماتریس مقایسات زوجی را مانند جدول (۲.۳) تشکیل دهد. ساعتی، قواعد کلی برای استفاده از AHP را به صورت زیر بیان می‌کند:

الف) قاعده‌ی شرایط دو طرفه: این قاعده بیان می‌کند که اگر معیار A، n مرتبه نسبت به معیار B اولویت داشته باشد، معیار B نیز $\frac{1}{n}$ ام مرتبه نسبت به معیار A اولویت دارد.

ب) قاعده‌ی تجانس: مقایسه‌ها فقط زمانی که عناصر قابل مقایسه باشند، معنی‌دار می‌باشد.



شکل ۳-۲: نمایش گرافیکی سلسله مراتب

جدول ۲.۳: ماتریس مقایسات زوجی برای تعیین وزن معیارها

هدف	D_1	D_2	D_3
D_1	۱		
D_2		۱	
D_3			۱

ج) قاعده‌ی استقلال: اهمیت نسبی (اولویت) یک عنصر در هر سطح به عناصری که در سطح پایین‌تر می‌باشند، بستگی ندارد.

د) قاعده‌ی انتظارات: این قاعده بیان می‌کند که سلسله مراتب باید به طور کامل آورده شود هیچ گزینه‌ای در سلسله مراتب، حذف یا به آن اضافه نمی‌شود.

در (AHP) عناصر هر سطح نسبت به عنصر مربوطه خود در سطح بالاتر به صورت دو به دو مقایسه شده و وزن آن‌ها محاسبه می‌گردد که این وزن‌ها را وزن نسبی می‌نامیم. سپس با تلفیق وزن‌های نسبی، وزن نهایی هرگزینه مشخص می‌گردد که آن را وزن مطلق می‌نامیم. برای محاسبه وزن هرگزینه از ماتریس

مقایسه زوجی، چندین روش پیشنهاد شده است که مهمترین آنها عبارتند از:

◇ روش حداقل مربعات^۶

◇ روش حداقل مربعات لگاریتمی^۷

◇ روش بردار ویژه^۸

◇ روش‌های تقریبی^۹

از آنجا که وزن شاخص‌ها منعکس کننده اهمیت آنها در تعیین هدف بوده و وزن هر گزینه نسبت به شاخص‌ها سهم آن گزینه در شاخص مربوطه می‌باشد می‌توان گفت که وزن نهایی هر گزینه از مجموع حاصلضرب وزن هر شاخص در وزن گزینه مربوط به آن شاخص بدست می‌آید. چنان که ملاحظه می‌شود در روش (AHP) وزن شاخص ارزیابی محاسبه شده توسط این روش، بر مبنای تجربه کارشناسان و قضاوت ذهنی وی است. به عبارت دیگر، چنانچه کارشناسان انتخاب شده تغییر کنند، وزن بدست آمده متفاوت خواهد بود و این یعنی اتکای صرف به قضاوت ذهنی تصمیم‌گیرنده که مهمترین نقص روش تحلیل سلسله مراتبی است [۴۴، ۵۳].

۳-۳ مراحل الگوریتم با رویکرد (AHP/DEA)

الگوریتم (AHP/DEA) یک الگوریتم دو مرحله‌ای برای وزن‌دهی کامل واحدهای تصمیم‌گیرنده دارای چندین ورودی و چندین خروجی است. در این الگوریتم ابتدا یک مدل تحلیل پوششی داده‌ها برای هر

^۶ Least squares method

^۷ Logarithmic least squares method

^۸ Eigenvector method

^۹ Approximation methods

زوج از واحدها بدون در نظر گرفتن سایر واحدها حل می‌گردد سپس با استفاده از نتایج بدست آمده از حل مدل‌های DEA به شیوه‌ای که گفته خواهد شد «ماتریس مقایسات زوجی» تشکیل و با حل مدل فرآیند تحلیل سلسله مراتبی یک سطحی وزن‌دهی انجام می‌شود (در این بخش هر رأس از شبکه را به عنوان یک واحد در نظر می‌گیریم). مراحل این الگوریتم را می‌توان به شکل زیر بیان کرد.

۳-۳-۱ مرحله‌ی اول: تشکیل ماتریس مقایسات زوجی با استفاده از DEA

فرض کنید n واحد وجود دارد. هر واحد I ورودی و J خروجی دارد. I_{ik} ورودی i -ام واحد k -ام و O_{jk} خروجی j -ام از واحد k -ام است. برای هر جفت واحد A و B یک مدل تحلیل پوششی داده‌ها به صورت زیر تشکیل می‌شود:

$$\begin{aligned} E_{AA} = \max \sum_j u_j O_{jA} \\ s.t \\ \sum_i v_i I_{iA} = 1 \\ \sum_j u_j O_{jA} \leq 1 \\ \sum_j u_j O_{jB} - \sum_i v_i I_{iB} \leq 0 \\ u_j, v_i \geq \epsilon \end{aligned} \quad (1-3)$$

فرض کنید E_{AA} مقدار بهینه‌ی کارایی واحد A باشد. جهت تضمین بهترین ارزیابی متقاطع برای واحد B مدل زیر طرح می‌شود:

$$\begin{aligned} E_{BA} = \max \sum_j u_j O_{jB} \\ s.t \\ \sum_i v_i I_{iB} = 1 \\ \sum_j u_j O_{jB} \leq 1 \\ \sum_j u_j O_{jA} - E_{AA} \sum_i v_i I_{iA} = 0 \\ u_j, v_i \geq \epsilon \end{aligned} \quad (2-3)$$

که E_{BA} مقدار بهینه‌ی ارزیابی متقاطع واحد B است. مشابه با دو موضوع بالا، دو مسئله AB و BB نیز بایستی حل شود تا E_{BB} و E_{AB} محاسبه گردد. به این ترتیب چهار مدل DEA حل و مقادیر E_{AA} ، E_{AB} ، E_{BB} و E_{BA} بدست می‌آیند. با بکارگیری نتایج مدل‌های فوق و با استفاده از رابطه زیر برای هر جفت واحدهای i, j ماتریس مقایسات زوجی که هر عنصر a_{ij} آن از رابطه زیر بدست می‌آید، تشکیل می‌گردد

$$a_{ij} = \begin{cases} 1 & i = j \\ \frac{E_{ii} + E_{ij}}{E_{jj} + E_{ji}} & i \neq j. \end{cases}$$

باید توجه کرد که در فرآیند تحلیل سلسله مراتبی قطراصلی ماتریس مقایسات زوجی دارای مقدار یک بوده و عنصر a_{ij} منعکس کننده ارزیابی واحد i نسبت به j است. در این ماتریس $a_{ji} = \frac{1}{a_{ij}}$ خواهد بود. ماتریس حاصل $n \times n$ است. این ماتریس، ارزیابی عینی تصمیم‌گیرنده را نشان می‌دهد که از طریق محاسبه مدل‌های زوجی تحلیل پوششی داده‌ها محاسبه شده و ماتریس ارزیابی متقاطعی را فراهم می‌سازد که اجازه می‌دهد هر واحد، مطلوب‌ترین ارزیابی را به نسبت واحد دیگر بدست آورد. لازم به ذکر است، چون آرایه‌های ماتریس مقایسات زوجی از طریق (DEA) حاصل شده‌اند و قضاوت کیفی در آن سهمی ندارد، پس عملاً ماتریس مقایسات زوجی حتماً سازگار خواهد بود.

فرض کنید ماتریس مقایسات زوجی حاصل از حل مدل‌های تحلیل پوششی داده‌ها را با A و درایه‌های آن را با a_{ij} نمایش دهیم. ماتریس A برای شبکه‌ای با n رأس بصورت زیر خواهد بود. که در آن برای

هر $i, a_{ii} = 1$ می‌باشد.

$$A = \begin{array}{c|ccccc} & 1 & 2 & \dots & n \\ \hline 1 & a_{11} & a_{12} & \dots & a_{1n} \\ 2 & a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ n & a_{n1} & a_{n2} & \dots & a_{nn} \end{array}$$

۳-۳-۲ مرحله‌ی دوم: رتبه‌بندی با استفاده از AHP.

در این مرحله برای تعیین وزن رؤوس (واحدها) از روش تقریبی میانگین حسابی استفاده می‌کنیم برای پیاده‌سازی این روش روی ماتریس مقایسات زوجی ابتدا ستون‌های ماتریس را نرمالیزه کرده و بعد از آن مجموع هر سطر را بدست آورده مقدار حاصل از نرمالیزه کردن بردار مجموع سطرها مقدار وزن واحد نسبت داده شده به آن سطر می‌باشد. به این ترتیب وزن هر رأس مشخص می‌گردد.

نویسنده بحث بالا را روی ماتریس A این‌گونه بیان می‌کند: ماتریس جدید A' که ستون‌های آن نرمالیزه شده از رابطه‌ی $a'_{ij} = \frac{a_{ij}}{\sum_{i=1}^n a_{ij}}$ بدست می‌آید. پس از بدست آوردن ماتریس A' وزن هر رأس از رابطه‌ی زیر بدست می‌آید.

$$a''_i = \sum_j a'_{ij} \quad \forall i \quad (\text{حاصل جمع هر سطر})$$

$$a'''_i = \frac{a''_i}{\sum_{i=1}^n a''_i} \quad (\text{وزن راس } i\text{-ام})$$

۳-۴ مزایای مدل تلفیقی تحلیل پوششی داده‌ها و تحلیل سلسله مراتبی

مزیت‌های استفاده از روش تلفیقی (AHP/DEA) نسبت به روش‌های قبلی (روش‌های موجود برای وزن‌دهی) به شرح زیر است:

(۱) نسبی بودن وزن‌ها (وزن‌ها در مقایسه یک رأس با دیگر رؤوس بدست می‌آیند).

(۲) کم یا زیاد کردن تعداد رؤوس، مقدار وزن‌ها را تغییر می‌دهد.

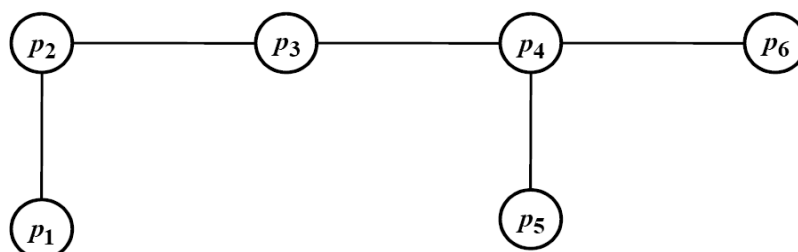
(۳) تعریف ورودی‌ها و خروجی‌ها برای رؤوس (واحدها) با توجه به نوع مسئله (اورژانس، آتش‌نشانی،

بیمارستان، مرکز پلیس ...) تغییر می‌کند (انعطاف‌پذیری زیاد الگوریتم).

- (۴) رابطه‌ی مستقیم وزن رؤوس با پارامترهای اساسی در تصمیم‌گیری مدیر برای انتخاب مکان.
- (۵) تلفیق روش کمی (تحلیل پوششی داده‌ها) و کیفی (فرآیند تحلیل سلسله مراتبی).
- (۶) دخیل بودن همزمان چند پارامتر در تعیین وزن رؤوس.
- (۷) هر چه وزن بدست آمده برای رأسی بزرگتر باشد یعنی آن رأس در برآورده کردن شاخص‌های مطلوب که مد نظر مدیر است موفق‌تر است.
- (۸) وزن دهی در این الگوریتم دارای روشی کاملاً ریاضی است و از هر گونه قضاوت‌های ذهنی بدور است.
- در مکانیابی وزن‌دار یکی از مراحل حساس و مؤثر، تعیین وزن رؤوس شبکه می‌باشد. بنابراین در این پایان‌نامه سعی بر آن شده است تا با استفاده از مدل تلفیقی تحلیل پوششی داده‌ها و تحلیل سلسله مراتبی این موضوع مهم به طور دقیق محاسبه شود.

۳-۵ مثال عددی

در این مثال ما ۶ شهر از استان سمنان را برای تعیین مکان بهینه فرودگاه مورد ارزیابی قرار می‌دهیم. در حقیقت شبکه‌ای خواهیم داشت که شهرها رؤوس و راه ارتباطی میان این شهرها یال‌های این شبکه خواهد بود. این شهرها عبارتند از: بیارجمند (p_1)، دامغان (p_2)، سمنان (p_4)، شاهرود (p_2)، گرمسار (p_6)، مهدیشهر (p_5). انتخاب ورودی و خروجی‌ها برای شهرها بر اساس نوع مسئله که همان تعیین مکان بهینه‌ی فرودگاه است و موجود بودن اطلاعات مزبور برای تمام شهرهای مورد بررسی و اینکه به نحوی توانایی کنترل و کم و زیاد کردن آنها وجود داشته باشد، صورت می‌گیرد. بر این اساس ورودی‌ها



شکل ۳-۳: نمودار شبکه‌ی شهرها

جدول ۳.۳: معرفی ورودی‌ها و خروجی‌ها

عنوان شاخص	نوع	مقیاس	نام پارامتر
جمعیت	خروجی	نفر	O_1
تعداد مراکز هواشناسی سینوپتیک	خروجی	عدد	O_2
فاصله از نزدیکترین فرودگاه	خروجی	کیلومتر	O_3
کیفیت راه ارتباطی	خروجی	دوقطبی*	O_4
جمعیت شاغلان غیر بومی غیر ساکن	خروجی	نفر	O_5
متوسط درآمد خانوار در سال	خروجی	ریال	O_6
متوسط تعداد جهانگردان و زائران به خارج از کشور در سال	خروجی	نفر	O_7
تعداد روزهای یخبندان در سال	ورودی	روز	I_1
متوسط بارندگی در سال	ورودی	میلیمتر	I_2
میانگین سرعت وزش باد	ورودی	متر بر ثانیه	I_3

* مقیاس دوقطبی عبارت است از تخصیص عددی بین $[0, 10]$ که عدد بزرگتر نشانه مطلوبیت بیشتر است

شامل تعداد روزهای یخبندان در سال، متوسط میزان بارندگی طی یک ماه و میانگین سرعت وزش باد و خروجی‌ها نیز شامل تعداد مراکز هواشناسی سینوپتیک، فاصله از نزدیکترین فرودگاه، کیفیت راه ارتباطی، جمعیت شاغلان غیر بومی غیر ساکن، متوسط درآمد خانوار در سال و متوسط تعداد زائران و جهانگردان به خارج از کشور در سال می‌باشد. نوع شاخص‌ها و مقیاس آنها و پارامترهای نسبت داده شده به هر یک در جدول (۳.۳) آمده است. مقدار شاخص‌ها (ورودی‌ها و خروجی‌ها) برای هریک از شهرها (رئوس) در جدول (۴.۳) مشخص شده است [۵۸]. در این سلسله مراتب هدف تعیین اهمیت شهرها

جدول ۴.۳: داده‌های ورودی و خروجی برای شهرها

	مهدیشهر	گرمسار	شاهرود	سمنان	دامغان	بیارجمند
O_1	۲۱۰۰۶	۳۹۵۲۳	۱۳۲۳۷۹	۱۲۶۷۸۰	۵۹۳۰۰	۲۵۰۴
O_2	۱	۱	۱	۱	۱	۱
O_3	۲۳۳	۱۰۰	۴۰۰	۲۱۶	۳۲۰	۵۱۸
O_4	۶	۷	۷	۸	۷	۴
O_5	۱۶۴	۳۰۲۱	۶۵۲۳	۷۲۸۴	۲۵۱۴	۸۹
O_6	۱۶۹۸۲۵۸۹	۱۸۳۹۸۵۴۱	۲۴۳۶۹۷۸۵	۲۵۳۹۸۷۳۲	۲۰۶۹۷۳۴۲	۱۴۴۴۴۲۴۰
O_7	۵۶۳	۱۱۳۰	۵۴۹۷	۳۵۱۴	۱۸۲۶	۱۲۹
I_1	۷۶	۴۵	۶۵	۵۵	۶۳	۸۲
I_2	۱۵/۲	۱۱/۲	۱۴	۱۴/۷	۹/۱	۱۳/۸
I_3	۱۳/۵	۱۶/۸	۱۴/۶	۱۳/۸	۲۱/۵	۱۲/۵

برای احداث فرودگاه می‌باشد؛ معیارهای تصمیم‌گیری در جدول (۳.۳) آمده است و گزینه‌ها شهرهای مهدیشهر، گرمسار، شاهرود، سمنان و دامغان و بیارجمند می‌باشند. حال با پیاده‌سازی روش ارائه شده برای تشکیل ماتریس مقایسات زوجی این ماتریس را پس از حل مدل‌های مربوطه به شکل جدول (۵.۳) داریم. فرض کنید می‌خواهیم درایه‌ی $a_{۱۴}$ در ماتریس مقایسات زوجی را بدست آوریم. اندیس یک مربوط به شهر بیارجمند و اندیس چهار برای شهر شاهرود می‌باشد. برای مقایسه‌ی این دو شهر باید مدل‌های $(۳-۳، ۳-۴، ۳-۵، ۳-۶)$ را حل کنیم. پس از تشکیل ماتریس مقایسات زوجی به کمک حل مدل‌های تحلیل پوششی داده‌ها برای تعیین وزن رئوس (وزن گزینه‌ها) از روش تقریبی میانگین حسابی استفاده می‌کنیم. وزن نهایی گزینه‌ها در جدول (۶.۳) آورده شده است. همانطور که در جدول (۶.۳) آمده است شاهرود بالاترین و مهدیشهر کمترین وزن را دارا هستند. و می‌توان پیش‌بینی کرد که

محل بهینه‌ی احداث فرودگاه به شاهرود نزدیک باشد.

$$E_{11} = \max[2504u_1 + u_2 + 518u_3 + 4u_4 + 89u_5 + 14444240u_6 + 129u_7]$$

s.t

$$82v_1 + 13/8v_2 + 12/5v_3 = 1$$

$$2504u_1 + u_2 + 518u_3 + 4u_4 + 89u_5 + 14444240u_6 + 129u_7 \leq 1$$

$$132379u_1 + u_2 + 400u_3 + 7u_4 + 6523u_5 + 24369785u_6 + 5497u_7$$

$$- [65v_1 + 14v_2 + 14/6v_3] \leq 0$$

$$u_1, \dots, u_7 \geq 0$$

$$v_1, v_2, v_3 \geq 0$$

(۳-۳)

$$E_{f1} = \max[132379u_1 + u_2 + 400u_3 + 7u_4 + 6523u_5 + 24369785u_6 + 5497u_7]$$

s.t

$$65v_1 + 14v_2 + 14/6v_3 = 1$$

$$132379u_1 + u_2 + 400u_3 + 7u_4 + 6523u_5 + 24369785u_6 + 5497u_7 \leq 1$$

$$2504u_1 + u_2 + 518u_3 + 4u_4 + 89u_5 + 14444240u_6 + 129u_7$$

$$- E_{11}[82v_1 + 13/8v_2 + 12/5v_3] = 0$$

$$u_1, \dots, u_7 \geq 0$$

$$v_1, v_2, v_3 \geq 0$$

(۴-۳)

$$\begin{aligned}
 E_{ff} = \max & [132379u_1 + u_2 + 400u_3 + 7u_4 + 6523u_5 + 24369785u_6 + 5497u_7] \\
 s.t \quad & 65v_1 + 14v_2 + 14/6v_3 = 1 \\
 & 132379u_1 + u_2 + 400u_3 + 7u_4 + 6523u_5 + 24369785u_6 + 5497u_7 \leq 1 \\
 & 2504u_1 + u_2 + 518u_3 + 4u_4 + 89u_5 + 14444240u_6 + 129u_7 \\
 & -[82v_1 + 13/8v_2 + 12/5v_3] \leq 0 \\
 & u_1, \dots, u_7 \geq 0 \\
 & v_1, v_2, v_3 \geq 0
 \end{aligned}$$

(۵-۳)

$$\begin{aligned}
 E_{1f} = \max & [2504u_1 + u_2 + 518u_3 + 4u_4 + 89u_5 + 14444240u_6 + 129u_7] \\
 s.t \quad & 82v_1 + 13/8v_2 + 12/5v_3 = 1 \\
 & 2504u_1 + u_2 + 518u_3 + 4u_4 + 89u_5 + 14444240u_6 + 129u_7 \leq 1 \\
 & 132379u_1 + u_2 + 400u_3 + 7u_4 + 6523u_5 + 24369785u_6 + 5497u_7 \\
 & -E_{ff}[65v_1 + 14v_2 + 14/6v_3] = 0 \\
 & u_1, \dots, u_7 \geq 0 \\
 & v_1, v_2, v_3 \geq 0
 \end{aligned}$$

(۶-۳)

پس از حل چهار مدل ارائه شده و به دست آوردن مقدار بهینه‌ی $E_{11}, E_{f1}, E_{ff}, E_{1f}$ درایه‌ی a_{1f} را به صورت زیر داریم

$$a_{1f} = \frac{E_{11} + E_{1f}}{E_{ff} + E_{f1}}.$$

در این مثال با توجه به تعداد گزینه‌ها و خواصی که برای ماتریس مقایسات زوجی گفته شد، نیاز است پانزده درایه‌ی این ماتریس را محاسبه کنیم که برای محاسبه‌ی هر درایه نیاز به حل چهار مدل برنامه‌ریزی

جدول ۵.۳: ماتریس مقایسات زوجی برای شهرها

مهدیشهر	گرمسار	شاهرود	سمنان	دامغان	بیارجمند
۱	۱	۱/۰۲۲۵	۱	۱	۱
۱	۱	۱	۱	۱	۱
۱	۱	۱	۱	۱	۱
۱/۰۲۸۷	۱	۱	۱	۱	۰/۹۷۸
۱	۱	۱	۱	۱	۱
۱	۱	۰/۹۲۷	۱	۱	۱

می‌باشد و در کل برای تکمیل کردن ماتریس مقایسات زوجی این مثال باید شصت مدل برنامه ریزی حل

کنیم.

جدول ۶.۳: وزن نهایی گزینه‌ها به روش میانگین حسابی

بیارجمند	۰/۱۶۷۳
دامغان	۰/۱۶۶۷
سمنان	۰/۱۶۶۷
شاهرود	۰/۱۶۸۲
گرمسار	۰/۱۶۶۷
مهدیشهر	۰/۱۶۶۴

فصل ۴

مدل تلفیقی مکانیابی-تحلیل پوششی داده‌ها

۴-۱ مدل تلفیقی مکانیابی-تحلیل پوششی داده‌های قطعی

سیستمی را در نظر بگیرید که در آن سرویس‌دهنده‌ها و سرویس‌گیرنده‌ها مشخص هستند و تعدادی نقاط به عنوان سرویس‌دهنده کاندیدای تخصیص به تعدادی سرویس‌گیرنده می‌باشند. هزینه‌ی استقرار و هزینه‌ی سرویس هر سرویس‌دهنده به هر یک از سرویس‌گیرنده‌ها از پیش تعیین شده است. حال برای محاسبه‌ی کارایی تخصیص، هر جفت (سرویس‌دهنده - سرویس‌گیرنده) را به عنوان یک واحد تصمیم‌گیری در نظر می‌گیریم و برای هر واحد شاخص‌هایی معرفی می‌کنیم.

هدف از تلفیق مسئله‌ی مراکز سرویس با ظرفیت سرویس محدود-نامحدود و تحلیل پوششی داده‌های قطعی، دستیابی به مجموعه‌ای از واحدها است که این مجموعه بیشترین مجموع کارایی و کمترین هزینه‌ی سرویس را دارا باشد. تلفیق به شکلی انجام می‌شود که سرویس‌دهنده‌ای که عمل سرویس‌دهی را انجام نمی‌دهد کارایی آن صفر در نظر گرفته شود. در حقیقت مدل پردازشگر را وادار می‌کند در صورتی که عمل سرویس‌دهی انجام نگیرد ضرایب شاخص‌ها برای آن واحد صفر در نظر گرفته شود. مدل حاصل از تلفیق، یک مدل دو هدفه است که تابع هدف اول آن ماکسیمم سازی مجموع کارایی تمام ترکیب‌های ممکن از سرویس‌دهنده‌ها و سرویس‌گیرنده‌ها است و تابع هدف دوم که خود از دو قسمت تشکیل شده است مینیمم سازی است. قسمت اول مینیمم سازی مجموع هزینه‌ی سرویس و قسمت دوم مینیمم سازی مجموع هزینه‌ی استقرار می‌باشد. مدل تلفیقی تخصیص مراکز سرویس با ظرفیت سرویس

نامحدود و تحلیل پوششی داده‌های قطعی به شکل زیر می‌باشد.

$$\begin{aligned}
 \max \quad & \sum_{k=1}^K \sum_{l=1}^L \sum_{j=1}^J u_{klj} O_{jkl} \\
 \min \quad & \sum_{k=1}^K \sum_{l=1}^L c_{kl} \text{dem}_l x_{kl} + \sum_{k=1}^K F_k y_k \\
 s.t \quad & \\
 & \sum_{k=1}^K x_{kl} = 1 \quad \forall l \\
 & \sum_{i=1}^I v_{kli} I_{ikl} = x_{kl} \quad \forall k, l \quad (a) \\
 & \sum_{j=1}^J u_{klj} O_{jrs} - \sum_{i=1}^I v_{kli} I_{irs} \leq 0 \quad \forall r, s, k, l \quad (1-4) \\
 & u_{klj} O_{jkl} \leq x_{kl} \quad \forall k, l, j \quad (b) \\
 & x_{kl} \leq y_k \quad \forall k, l \\
 & u_{klj} \geq \epsilon x_{kl} \quad , \quad v_{kli} \geq \epsilon x_{kl} \quad \forall k, l, i, j \quad (c) \\
 & u_{klj}, v_{kli} \geq 0 \quad \& \quad x_{kl}, y_k \in \{0, 1\}
 \end{aligned}$$

که متغیرها و پارامترهای مدل بصورت زیر می‌باشند.

u_{klj} : وزن خروجی j -ام واحد kl .

O_{jkl} : خروجی j -ام واحد kl .

v_{kli} : وزن ورودی i -ام واحد kl .

I_{ikl} : ورودی i -ام واحد kl .

C_{kl} : هزینه‌ی هر واحد سرویس از سرویس‌دهنده‌ی k به سرویس‌گیرنده‌ی l .

dem_l : میزان سرویس مورد نیاز برای سرویس‌گیرنده‌ی l .

F_k : هزینه‌ی ثابت استقرار (بازبودن) واحد سرویس‌دهنده‌ی k .

ϵ : یک مقدار کوچک و مثبت است برای ممانعت از صفر شدن وزن شاخص‌ها.

x_{kl}, y_k متغیرهای تصمیم مسأله هستند که در فصل ۱ معرفی شده‌اند.

وقتی که سرویس‌دهنده‌ی k به سرویس‌گیرنده‌ی l سرویس می‌دهد ($x_{kl} = 1$)، همانطور که قبلاً

گفته شده است برای منطقی‌تر شدن جواب نیاز است که وزن تخصیص یافته به شاخص‌های مربوط

به واحدهایی که عمل سرویس دهی را انجام می‌دهند بزرگتر از صفر گردد که این موضوع در محدودیت (c) آورده شده است. وقتی که سرویس‌دهنده‌ی k به سرویس‌گیرنده‌ی l سرویس ندهد ($x_{kl} = 0$)، محدودیت (c) به پردازشگر این اجازه را می‌دهد که برای شاخص‌های این نوع واحدها وزن صفر را هم در نظر بگیرد و در مرحله‌ی بعد محدودیت‌های (a, b) پردازشگر را وادار می‌کنند که در صورت عدم سرویس‌دهی، وزن شاخص‌ها برای واحد مربوطه صفر گردد. با توجه به تعریف یک واحد تصمیم‌گیری در این مدل چون برای هر جفت (سرویس‌دهنده-سرویس‌گیرنده) شاخص‌ها بصورت جداگانه مشخص شده‌است پس هر ترکیبی از سرویس‌دهنده و سرویس‌گیرنده‌ها کارایی مربوط به خود را دارا است و مدل سعی می‌کند ترکیبی را انتخاب کند که بیشترین کارایی را داشته باشد [۲۲].

دو مدل تخصیص مراکز سرویس با ظرفیت سرویس محدود و تحلیل پوششی داده‌های قطعی را با این رویکرد که کارایی واحدی که عمل سرویس‌دهی را انجام نمی‌دهد در تابع هدف اول محاسبه نگردد (صفر در نظر گرفته شود) تلفیق می‌نماییم. برای این منظور مدل تلفیقی بگونه‌ای طراحی می‌گردد که وزن شاخص‌های مربوط به واحدی که عمل سرویس‌دهی را انجام نمی‌دهد صفر در نظر گرفته می‌شود. مدل تلفیقی تخصیص مراکز سرویس با ظرفیت سرویس محدود و تحلیل پوششی داده‌های قطعی بصورت

زیر است.

$$\begin{aligned}
 \max \quad & \sum_{k=1}^K \sum_{l=1}^L \sum_{j=1}^J u_{klj} O_{jkl} \\
 \min \quad & \sum_{k=1}^K \sum_{l=1}^L c_{kl} b_{kl} x_{kl} + \sum_{k=1}^K F_k y_k \\
 s.t \quad & \\
 & \sum_{k=1}^K x_{kl} \geq 1 \quad \forall l \\
 & \sum_{i=1}^I v_{kli} I_{ikl} = x_{kl} \quad \forall k, l \\
 & \sum_{j=1}^J u_{klj} O_{jrs} - \sum_{i=1}^I v_{kli} I_{irs} \leq 0 \quad \forall r, s, k, l \\
 & \sum_{k=1}^K b_{kl} = dem_l \quad \forall l \quad (2-4) \\
 & b_{kl} \leq \min[dem_l, O_{mkl}] y_k \quad \forall k, l \quad (a') \\
 & x_{kl} \leq y_k \quad \forall k, l \\
 & b_{kl} \geq x_{kl} \quad \forall k, l \quad (b') \\
 & u_{klj} O_{jkl} \leq x_{kl} \quad \forall k, l, j \\
 & u_{klj} \geq \epsilon x_{kl} \quad , \quad v_{kli} \geq \epsilon x_{kl} \quad \forall k, l, i, j \\
 & u_{klj}, v_{kli} \geq 0 \quad \& \quad x_{kl}, y_k \in \{0, 1\}
 \end{aligned}$$

با توجه به اینکه هر سرویس دهنده‌ای که ظرفیت سرویس بالاتری داشته باشد انتخاب آن برای تخصیص مفیدتر (به صرفه‌تر) خواهد بود، و هر چه خروجی یک واحد بزرگتر باشد آن واحد بهره‌وری (کارایی) بیشتری خواهد داشت ظرفیت سرویس هر سرویس دهنده در زمان اتصال به هر یک از سرویس گیرنده‌ها به عنوان یک خروجی در نظر گرفته می‌شود. در محدودیت (a') ظرفیت سرویس هر سرویس دهنده در زمان اتصال به هر یک از سرویس گیرنده‌ها به عنوان یک خروجی با پارامتر (O_{mkl}) مشخص شده است. محدودیت (b') به شکلی طراحی شده است که اگر واحدی عمل سرویس دهی را انجام دهد $(x_l = 1)$ باید حداقل یک واحد سرویس را از سرویس دهنده‌ی k به سرویس گیرنده‌ی l منتقل کند یا اگر سرویسی از سرویس دهنده‌ی k به سرویس گیرنده‌ی l منتقل نگردد $(b_{kl} = 0)$ آنگاه متغیر تصمیم x_{kl} برابر صفر می‌گردد و کارایی واحد kl در تابع هدف محاسبه نمی‌گردد (صفر در نظر گرفته می‌شود) [۲۲].

۴-۱-۱ معرفی روش وزن دهی به اهداف برای حل مسائل چندهدفه

در شرایط واقعی وضعیت‌هایی رخ می‌دهد که مدلی ریاضی شامل تنها یک هدف، بیان‌گر واقعیت و خواسته‌های مورد نظر تصمیم‌گیرنده نبوده و این امر کارایی و مطلوبیت نتایج حاصل از مدل را کاهش می‌دهد. مسائل بسیاری در عالم واقع می‌توان یافت که تحت تأثیر چند هدف متعدد قرار می‌گیرند. مثلاً در هر کارخانه اهداف متعدد و بعضاً متناقضی برای قسمت‌های مختلف می‌توان یافت. از جمله در قسمت تولید اهدافی مانند: حداکثر کردن درآمد، حداقل کردن هزینه، حداکثر کردن استفاده از ظرفیت تولید و حداقل کردن میزان موجودی، در قسمت نگهداری و تعمیرات اهدافی مانند: حداکثر کردن زمان استفاده از ماشین آلات و حداقل کردن موجودی قطعات یدکی و در قسمت فروش اهدافی مانند: تنوع محصولات، حداقل کردن هزینه‌های فروش. بعید نیست هدف هر قسمت با اهداف قسمت‌های دیگر متفاوت و گاه متضاد باشد. شکل ریاضی مسائل با اهداف چندگانه به صورت زیر است:

$$\begin{aligned} \max(\min) \quad & Z = [z_1, z_2, \dots, z_p] \\ & z_1 = z_1(x_j) \\ & \vdots \\ & z_{p-1} = z_{p-1}(x_j) \\ & z_p = z_p(x_j) \end{aligned} \quad (3-4)$$

s.t

$$\begin{aligned} g_i(x_j) &\leq b_i \quad i = 1, 2, \dots, m \\ x_j &\geq 0 \quad j = 1, 2, \dots, n \end{aligned}$$

که $z_p(x_j)$ و $g_i(x_j)$ توابعی خطی از متغیر تصمیم x_j و b_i مقداری نامنفی و ثابت است. مدل فوق دارای p تابع هدف، m محدودیت و n متغیر است. جواب ایده‌آل برای این مدل، مجموعه مقادیر اختصاص یافته به متغیرهای تصمیم x_j است که در تمامی m محدودیت صدق کرده و هم زمان p تابع هدف را حداکثر (حداقل) کند. با این همه جوابی که بتواند یک هدف را بهینه کند به علت وجود اهداف متضاد

در مدل ممکن است نتواند سایر اهداف را بهینه نماید. لذا در اینگونه مسائل، بهینه شدن یک تابع هدف لزوماً مترادف بهینه شدن تمامی توابع هدف و مسأله نیست. ما در این قسمت به معرفی روش وزن‌دهی به اهداف برای مسائل چند هدفه می‌پردازیم.

در این روش تصمیم‌گیرنده به اهداف مختلف وزن داده و از ترکیب اهداف وزن داده شده یک تابع هدف به دست می‌آورد. به این ترتیب، مدلی با یک تابع هدف ایجاد می‌شود که مانند مسائل برنامه‌ریزی خطی قابل حل خواهد بود. وزن هر کدام از اهداف مقادیری نامنفی فرض و مجموع وزن‌های داده شده برابر یک در نظر گرفته می‌شود. به این ترتیب، هر گاه در مدلی اهمیت هدف اول چهار برابر هدف دوم باشد، برای این کار کافی است مجموع اهمیت هر دو هدف $(4+1=5)$ را پیدا کرده، وزن هدف اول را از تقسیم اهمیت آن هدف بر مجموع حاصل یعنی $\frac{4}{5} = 0.8$ به دست آورد. پس، وزن هدف دوم برابر $\frac{1}{5} = 0.2$ خواهد شد. لذا تابع هدف مسأله چنین است

$$\max(\min) Z = 0.8(\text{تابع هدف اول}) + 0.2(\text{تابع هدف دوم})$$

گام‌های لازم برای وزن دادن به اهداف بصورت زیر است [۵۷].

گام ۱: اهمیت هر هدف را نسبت به هدف یا اهداف دیگر مشخص و وزن‌ها را تعیین کنید.

گام ۲: تمامی توابع هدف را باید به صورت max و یا به صورت min در آوردید.

گام ۳: مطمئن شوید که ضرایب هر تابع هدف نسبت به دیگر توابع هدف دارای یک درجه‌ی بزرگی می‌باشد.

گام ۴: با ضرب کردن وزن‌ها در توابع هدف یک تابع هدف تکی ایجاد کنید.

گام ۵: مسأله برنامه‌ریزی خطی ناشی از گام ۴ را حل کنید.

گام ۶: جواب بهینه‌ی به دست آمده از حل مسئله گام ۵ را در هر کدام از توابع هدف قرار داده و z_i^* ها را به دست آورید.

۴-۱-۲ مثال عددی

می‌خواهیم مدل تلفیقی تخصیص مراکز سرویس با ظرفیت سرویس محدود-نامحدود و تحلیل پوششی داده‌های قطعی را برای داده‌های فرضی زیرپایه سازی کنیم.

در این مثال ما سیستمی را در نظر می‌گیریم که شامل دو سرویس‌دهنده و چهار سرویس‌گیرنده می‌باشد. از آنجا که هر زوج (سرویس‌دهنده-سرویس‌گیرنده) را به عنوان یک واحد تصمیم‌گیری در نظر گرفته‌ایم پس هشت واحد تصمیم‌گیری خواهیم داشت. برای هر واحد سه ورودی و دو خروجی و $\epsilon = 10^{-4}$ در نظر می‌گیریم. مقدار پارامترهای مربوط به ورودی و خروجی واحدها و هزینه‌ی هر واحد سرویس در جدول (۱.۴)، میزان سرویس مورد نیاز برای هر سرویس‌گیرنده و هزینه‌ی استقرار مراکز برای هر سرویس‌دهنده در جدول (۲.۴) آمده است (F_k معرف سرویس‌دهنده‌ی k -ام و D_l معرف سرویس‌گیرنده‌ی l -ام است). میزان تغییراندیس‌های بکاررفته در مدل (۴-۱) و (۴-۲) برای این مثال به شرح زیر است.

$$k = 1, 2 \quad l = 1, 2, 3, 4 \quad i = 1, 2, 3$$

$$r = 1, 2 \quad s = 1, 2, 3, 4 \quad j = 1, 2$$

برای حالت ظرفیت محدود همانطور که گفته شد ظرفیت سرویس به عنوان یک خروجی در نظر گرفته می‌شود پس در این حالت سیستم شامل سه خروجی خواهد بود که خروجی سوم همان ظرفیت سرویس خواهد بود. حال با در نظر گرفتن داده‌های جدولی ارائه شده مدل تخصیص مراکز با ظرفیت سرویس نامحدود و تحلیل پوششی داده‌های قطعی را پیاده‌سازی می‌کنیم. نتایج حاصل از حل مدل تلفیقی به

جدول ۱.۴: داده‌های مربوط به شاخص‌ها

سرویس دهنده	سرویس گیرنده	هزینه‌ی هر واحد سرویس	I_1	I_2	I_3	O_1	O_2
F_1	D_1	۴۶	۷۶	۶۹	۶۳	۷۳	۴
F_1	D_2	۶۹	۵۰	۹۲	۷۵	۸۰	۳۵
F_1	D_3	۶۲	۵۰	۷۶	۶۴	۶۹	۳۷
F_1	D_4	۳۹	۵۰	۲۵	۲۴	۲۱	۳۶
F_2	D_1	۱۱	۹۸	۶۹	۴۱	۶۳	۳۴
F_2	D_2	۷۲	۲۹	۵۹	۷۵	۱۱	۸۲
F_2	D_3	۷۴	۴۰	۳۹	۱۸	۴۹	۶۰
F_2	D_4	۷۸	۵۶	۴۲	۹۰	۵۸	۸۲

جدول ۲.۴: هزینه‌ی استقرار و میزان سرویس مورد نیاز

میزان سرویس مورد نیاز هر سرویس گیرنده		هزینه‌ی استقرار هر سرویس دهنده			
F_1	F_2	D_1	D_2	D_3	D_4
۲۵۰	۳۷۰	۸	۲۸	۲۰	۱۷

روش وزن‌دهی به اهداف با در نظر گرفتن وزن α برای تابع هدف کارایی و وزن β برای تابع هدف هزینه به شرح جدول (۳.۴) و (۴.۴) می‌باشد. نمایش گرافیکی تخصیص با ظرفیت سرویس نامحدود و با رویکرد هزینه و کارایی در شکل (۴-۱) نمایش داده شده است. حال با توجه به اینکه برای حل مدل تلفیقی با ظرفیت محدود نیاز به میزان ظرفیت سرویس هر سرویس دهنده داریم مقدار این پارامتر در جدول (۵.۴) آورده شده است. حال با در نظر گرفتن داده‌های جدولی ارائه شده مدل تخصیص مراکز با ظرفیت سرویس محدود و تحلیل پوششی داده‌های قطعی را پیاده‌سازی می‌کنیم. نتایج حاصل از حل مدل تلفیقی به روش وزن‌دهی به اهداف با در نظر گرفتن وزن α برای تابع هدف کارایی و وزن β برای تابع هدف هزینه به شرح جدول (۶.۴) و (۷.۴) می‌باشد. نمایش گرافیکی تخصیص با ظرفیت سرویس محدود و با رویکرد هزینه و کارایی در شکل (۴-۲) نمایش داده شده است.

جدول ۷.۴: نتایج حل مثال با ظرفیت سرویس محدود با در نظر گرفتن $\alpha = 1, \beta = 0$

متغیرهای تصمیم													تابع هدف		
b_{11}	b_{12}	b_{13}	b_{14}	b_{22}	b_{23}	b_{24}	x_{11}	x_{14}	x_{23}	x_{24}	y_1	y_2	کل	کارایی	هزینه
۸	۲۰	۱۹	۱۶	۸	۱	۱	۱	۱	۱	۱	۱	۱	۵/۸۵	۵/۸۵	۴۹۰۷

۲-۴ مدل تلفیقی مکانیابی - تحلیل پوششی داده‌های بازه‌ای

در این بخش در نظر داریم مدل تخصیص مراکز سرویس و تحلیل پوششی داده‌ها را با شاخص‌های (ورودی و خروجی‌های) بازه‌ای تلفیق نمائیم. در این تلفیق از رویکردی که در تلفیق اولیه آورده شده است استفاده می‌کنیم. در حقیقت همان مدل بخش (۴-۱) را داریم که ورودی و خروجی‌ها برای آن بازه‌ای شده‌اند. پس از خطی کردن مدل بازه‌ای با همان روندی که در معرفی تحلیل پوششی داده‌های بازه‌ای گفته شد مدل تلفیقی تخصیص مراکز سرویس با ظرفیت سرویس نامحدود و تحلیل پوششی داده‌های بازه‌ای را به شکل زیر خواهیم داشت.

$$\begin{aligned}
 & \max \quad \sum_k \sum_l \sum_j [u_{klj} \underline{Q}_{jkl} + p_{jkl} (\overline{Q}_{jkl} - \underline{Q}_{jkl})] \\
 & \min \quad \sum_k \sum_l c_{kl} \text{dem}_l x_{kl} + \sum_k F_k y_k \\
 & s.t. \\
 & \sum_j [u_{klj} \underline{Q}_{jrs} + p_{jrs} (\overline{Q}_{jrs} - \underline{Q}_{jrs})] - \sum_i [v_{kli} \underline{I}_{irs} + q_{irs} (\overline{I}_{irs} - \underline{I}_{irs})] \leq 0 \\
 & \sum_i [v_{kli} \underline{I}_{ikl} + q_{ikl} (\overline{I}_{ikl} - \underline{I}_{ikl})] = x_{kl} \quad \forall k, l, \quad \sum_k x_{kl} = 1 \\
 & u_{klj} \underline{Q}_{jkl} + p_{jkl} (\overline{Q}_{jkl} - \underline{Q}_{jkl}) \leq x_{kl}, \quad u_{klj}, v_{kli} \geq \epsilon x_{kl} \\
 & p_{jkl} - u_{klj} \leq 0, \quad q_{ikl} - v_{kli} \leq 0 \\
 & p_{jkl}, q_{ikl}, u_{klj}, v_{kli} \geq 0, \quad x_{kl}, y_k \in \{0, 1\}
 \end{aligned} \tag{۴-۴}$$

مدل تلفیقی تخصیص مراکز سرویس با ظرفیت محدود و تحلیل پوششی داده‌های بازه‌ای به شکل زیر خواهد بود:

$$\begin{aligned}
 & \max \quad \sum_k \sum_l \sum_j [u_{klj} \underline{Q}_{jkl} + p_{jkl} (\overline{O}_{jkl} - \underline{Q}_{jkl})] \\
 & \min \quad \sum_k \sum_l c_{kl} b_{kl} + \sum_k F_k y_k \\
 & s.t \quad \sum_j [u_{klj} \underline{Q}_{jrs} + p_{jrs} (\overline{O}_{jrs} - \underline{Q}_{jrs})] - \sum_i [v_{kli} \underline{I}_{irs} + q_{irs} (\overline{I}_{irs} - \underline{I}_{irs})] \leq 0 \\
 & b_{kl} \leq \min[dem_l, cap_k] y_k, \quad \sum_k b_{kl} = dem_l \quad (5-4) \\
 & \sum_i [v_{kli} \underline{I}_{ikl} + q_{ikl} (\overline{I}_{ikl} - \underline{I}_{ikl})] = x_{kl} \quad \forall k, l, \quad \sum_k x_{kl} \geq 1 \\
 & u_{klj} \underline{Q}_{jkl} + p_{jkl} (\overline{O}_{jkl} - \underline{Q}_{jkl}) \leq x_{kl}, \quad u_{klj}, v_{kli} \geq \epsilon x_{kl} \\
 & p_{jkl} - u_{klj} \leq 0, \quad q_{ikl} - v_{kli} \leq 0 \\
 & p_{jkl}, q_{ikl}, u_{klj}, v_{kli} \geq 0, \quad x_{kl}, y_k \in \{0, 1\}
 \end{aligned}$$

پارامتر cap_k را ظرفیت سرویس هر سرویس‌دهنده تعریف می‌کنیم. ما در این مدل ظرفیت سرویس هر سرویس‌دهنده را دیگر بعنوان یک خروجی تعریف نمی‌کنیم بلکه بعنوان یک پارامتر مشخص برای هر سرویس‌دهنده می‌شناسیم. این مدل‌ها را با در نظر گرفتن مقدار قطعی برای بعضی از ورودی‌ها و خروجی‌ها و مقدار بازه‌ای برای دیگر ورودی‌ها و خروجی‌ها می‌توان حل کرد. در مدل‌های تلفیقی موجود برای مسئله مکانیابی-تحلیل پوششی داده‌ها، مدیر (تصمیم‌گیرنده) برای تعیین میزان ورودی و خروجی‌ها مجبور به تعیین مقدار قطعی است. در حالی که در عالم واقعیت در بعضی از مواقع تعیین مقدار قطعی برای پارامترها، کاری سخت است. حال با مدل ارائه شده در این بخش این اختیار به تصمیم‌گیرنده داده شده است که شاخص‌های ارزیابی کارائی واحدهای خود را بصورت بازه‌ای تعیین کند. کار با این مدل‌ها با توجه به شرایط جدیدی که در آنها وارد شده برای تصمیم‌گیرنده راحت‌تر است. با توجه به اینکه مسائل موجود در عالم واقعیت به این مدل‌ها نزدیک‌تر بوده پس می‌توان گفت برای، پیاده‌سازی مسائل واقعی این مدل کاراتر و کاربردی‌تر است.

۳-۴ مدل تلفیقی مکانیابی-تحلیل پوششی داده‌های تصادفی

در این بخش با توجه به اینکه در مسائل دنیای واقعی امکان دارد ورودی و خروجی‌هایی که برای ارزیابی کارایی مورد استفاده قرار می‌گیرد ماهیت تصادفی داشته باشند، حالت تصادفی مسئله مکانیابی-تحلیل پوششی داده‌ها را تحلیل و برای آن مدلی ارائه می‌کنیم. همانطور که در قسمت (۲-۵) گفته شد شکل‌های مختلفی از مدل‌های تحلیل پوششی داده‌های تصادفی تا به حال مطرح شده است که ما در این زیربخش مدل تحلیل پوششی داده‌های آینده‌نگر را برای تلفیق با مسئله تخصیص مراکز سرویس با ظرفیت نامحدود انتخاب کرده‌ایم. در مدل تلفیقی ارائه شده در این بخش ما مقدار خروجی را بصورت تصادفی در نظر گرفته‌ایم. در این مدل سعی شده میزان قبول خطاهای تصادفی در محاسبه کارایی واحدهای تصمیم‌گیرنده را به میزان رضایت‌مندی نتایج مدل برای تصمیم‌گیرنده مربوط سازیم. طبیعی است که این مدل برای سیستم‌هایی که ورودی‌های مشخص (قطعی) و خروجی‌های تصادفی دارند مانند: شعب بانک‌ها یا بیمارستان‌ها و غیره کاربرد دارد. حال مدل تلفیقی تخصیص مراکز سرویس با ظرفیت نامحدود و تحلیل پوششی داده‌های آینده‌نگر را به شکل زیر ارائه می‌دهیم:

$$\begin{aligned}
 & \max \quad \sum_k \sum_l \sum_j u_{klj} \bar{O}_{jkl} \\
 & \min \quad \sum_k \sum_l c_{kl} \text{dem}_l x_{kl} + \sum_k F_k y_k \\
 & s, t \\
 & \sum_k x_{kl} \geq 1 \\
 & x_{kl} = y_k \\
 & \sum_i v_{kli} I_{ikl} = x_{kl} \\
 & \beta_{rs} \sum_i v_{kli} I_{irs} - \sum_j u_{klj} \bar{O}_{jrs} \geq (\sum_j u_{klj} b_{jrs}) F^{-1} (1 - \alpha_{rs}) \\
 & u_{klj} \bar{O}_{jkl} \leq x_{kl} \quad (i) \\
 & u_{klj} \geq \epsilon x_{kl} \quad \& \quad v_{kli} \geq \epsilon x_{kl} \\
 & u_{klj}, v_{kli} \geq 0 \quad \& \quad x_{kl}, y_k \in \{0, 1\}
 \end{aligned}
 \tag{۶-۴}$$

در محدودیت (i) می‌خواهیم این محدودیت را طوری مطرح کنیم که در صورت صفر شدن x_{kl} (یعنی سرویس‌دهنده k به سرویس‌گیرنده l سرویس ندهد) وزن خروجی‌های مربوط به این واحد (واحد kl) صفر گردد از طرفی می‌خواهیم ساده‌سازی به شکلی صورت گیرد که متغیر تصادفی $\hat{O}_{jkl}$ جای خود را به مقدار موردانتظار خروجی z -ام واحد سرویس‌دهنده k که به سرویس‌گیرنده l سرویس می‌دهد ($\bar{O}_{jkl}$)، بدهد. برای رسیدن به خواسته خود از این نکته استفاده می‌کنیم که امید ریاضی متغیر تصادفی نامنفی زمانی صفر می‌گردد که آن متغیر صفر باشد. همان طور که در قسمت (۲-۵) گفته شد متغیر تصادفی $\hat{O}_{jkl}$ را می‌توانیم به صورت زیر داشته باشیم:

$$\hat{O}_{jkl} = \bar{O}_{jkl} \pm b_{jkl}\zeta$$

حال با توجه به اینکه ζ یک متغیر تصادفی نرمال استاندارد است نتایج زیر را خواهیم داشت:

$$E(u_{klj}\hat{O}_{jkl}) = u_{klj}\bar{O}_{jkl}$$

$$if \quad u_{klj}\bar{O}_{jkl} = 0 \xrightarrow{\text{چون } \bar{O}_{jkl} > 0 \text{ و } u_{klj} \geq 0} u_{klj} = 0$$

پس همانطور که در مدل بالا آورده شده است. محدودیت مورد نظر را به شکل زیر داریم:

$$E(u_{klj}\hat{O}_{jkl}) \leq x_{kl} \implies u_{klj}\bar{O}_{jkl} \leq x_{kl}$$

مزیت بکارگیری حالت تصادفی برای تلفیق، نسبت به روش تحلیل پوششی داده‌های قطعی این است که مدیریت می‌تواند با قبول ریسک متفاوت کارایی واحدهای خود را پیش‌بینی نماید و به این ترتیب با در نظر داشتن وضعیت آینده‌ی واحدها جهت بهبود عملکرد واحدهای ناکارا اقدام کند. مجموعه واحدهای منتخب که از حل مدل تلفیقی بدست می‌آیند بالاترین کارایی و کمترین میزان هزینه را دارا هستند.

۴-۳-۱ مثال عددی

سیستمی را در نظر می‌گیریم که شامل دو سرویس‌دهنده و چهار سرویس‌گیرنده می‌باشد. از آنجا که هر زوج (سرویس‌دهنده-سرویس‌گیرنده) را به عنوان یک واحد تصمیم‌گیری در نظر گرفته‌ایم پس هشت واحد تصمیم‌گیری خواهیم داشت. برای هر واحد سه ورودی و دو خروجی و $\epsilon = 10^{-4}$ در نظر می‌گیریم. خروجی‌های هر واحد را نیز تصادفی در نظر گرفته و مقدار مورد انتظار و انحراف معیار درونی برای هر خروجی در جدول (۸.۴) آمده است.

مقدار کارایی مورد انتظار برای هر واحد را یک در نظر می‌گیریم ($\beta_s = 1 \quad \forall s$). میزان ریسک‌پذیری تصمیم‌گیرنده را پنج صدم در نظر می‌گیریم ($\alpha_s = 0.5, \quad \forall s$) و یا عبارتی احتمال رسیدن به سطح مطلوب کارایی برای تمام واحدها (۰/۹۵) در نظر گرفته می‌شود. مقدار پارامترهای مربوط به ورودی و خروجی واحدها و هزینه‌ی هر واحد سرویس در جدول (۱۰.۴)، میزان سرویس مورد نیاز برای هر سرویس‌گیرنده و هزینه‌ی استقرار مراکز برای هر سرویس‌دهنده در جدول (۲.۴) آمده است (F_k معرف سرویس‌دهنده‌ی k -ام و D_l معرف سرویس‌گیرنده‌ی l -ام است). حال با در نظر گرفتن داده‌های جدولی ارائه شده مدل تلفیقی مکانیابی-تحلیل پوششی داده‌های تصادفی را پیاده‌سازی می‌کنیم. نتایج حاصل از حل مدل تلفیقی به روش وزن‌دهی به اهداف با در نظر گرفتن وزن α برای تابع هدف کارایی و وزن β برای تابع هدف هزینه به شرح جدول (۹.۴) و (۱۰.۴) می‌باشد. نمایش گرافیکی تخصیص در حالت تصادفی و با رویکرد هزینه و کارایی در شکل (۳-۴) نمایش داده شده است.

جدول ۸.۴: داده‌های مربوط به شاخص‌های تصادفی

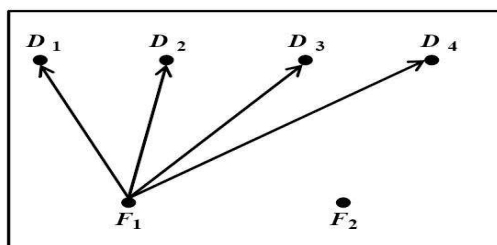
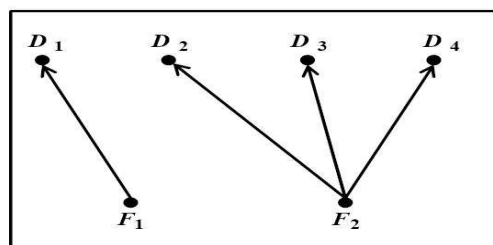
سرویس دهنده	سرویس گیرنده	مقدار منتظره		انحراف معیار	
		خروجی اول	خروجی دوم	خروجی اول	خروجی دوم
F_1	D_1	۷۲/۸۳	۳/۸۳	۰/۹۱	۰/۷۱
F_1	D_2	۷۹/۸۳	۳۴/۸۳	۰/۷۱	۰/۹۱
F_1	D_3	۶۸/۵	۳۶/۸۳	۰/۹۱	۰/۷۱
F_1	D_4	۲۱	۳۵/۶۶	۰/۵۸	۱
F_2	D_1	۶۸/۶۶	۳۴/۱۶	۰/۸۲	۰/۹۱
F_2	D_2	۱۰/۵	۸۲	۰/۹۱	۰/۸۲
F_2	D_3	۴۵/۶۶	۵۹/۸۳	۰/۸۲	۰/۹۱
F_2	D_4	۵۷/۸۳	۸۲/۱۶	۰/۹۱	۰/۹۱

جدول ۹.۴: نتایج حل مثال (۱-۳-۴) با در نظر گرفتن $\alpha = 0, \beta = 1$

x_{11}	x_{12}	x_{13}	x_{14}	y_1	y_2	تابع هدف هزینه	تابع هدف کارایی	تابع هدف کل
۱	۱	۱	۱	۱	۰	۴۴۵۳	۳/۸۷	-۴۴۵۳

جدول ۱۰.۴: نتایج حل مثال (۱-۳-۴) با در نظر گرفتن $\alpha = 1, \beta = 0$

x_{11}	x_{22}	x_{23}	x_{24}	y_1	y_2	تابع هدف هزینه	تابع هدف کارایی	تابع هدف کل
۱	۱	۱	۱	۱	۱	۵۹۶۳	۴	۴

 $\alpha = 0$, $\beta = 1$  $\alpha = 1$, $\beta = 0$

شکل ۳-۴: نمایش گرافیکی تخصیص در حالت تصادفی با رویکرد هزینه و کارایی

۴-۴ مدل تلفیقی مسئله مکانیابی-DEA با در نظر گرفتن وزن نسبی شاخص‌ها

با توجه به این که با در نظر گرفتن $\alpha = 1$ مدل ارائه شده در رویکرد دوم (بخش ۲-۶) همان مدل رویکرد اول است بنابراین مدل تلفیقی را برای رویکرد دوم که بنظر جامع‌تر است بیان می‌کنیم. در حقیقت مدل ارائه شده در این بخش همان مدل ارائه شده در بخش (۴-۱) است که نظر کارشناسان از طریق «محدودیت‌های کنترل وزن نسبی» به آخر افزوده شده است. مزیت‌های افزوده شدن محدودیت‌های کنترل وزن نسبی به مدل قبلاً در قسمت (۲-۶) آورده شده است. مدل تلفیقی مسئله تخصیص مراکز سرویس با ظرفیت نامحدود و تحلیل پوششی داده‌ها با در نظر گرفتن وزن نسبی شاخص‌ها به شکل زیر خواهد بود:

$$\begin{aligned}
 & \max \quad \sum_k \sum_l \sum_j u_{klj} O_{jkl} \\
 & \min \quad \sum_k \sum_l c_{kl} \text{dem}_l x_{kl} + \sum_k F_k y_k \\
 & s.t. \\
 & u_{klj} O_{jkl} \leq x_{kl} \quad , \quad x_{kl} \leq y_k \\
 & \sum_j u_{klj} O_{jrs} - \sum_i v_{kli} I_{irs} \leq 0 \quad , \quad \sum_k x_{kl} = 1 \quad (۷-۴) \\
 & (w_{kli} \alpha \sum_i v_{kli}) - v_{kli} \leq 0 \quad , \quad \sum_i v_{kli} I_{ikl} = x_{kl} \\
 & v_{kli} - x_{kl} (1 - \alpha + w_{kli} \alpha) \sum_i v_{kli} \leq 0 \quad , \quad (w_{klj} \alpha \sum_j u_{klj}) - u_{klj} \leq 0 \\
 & u_{klj} - x_{kl} (1 - \alpha + w_{klj} \alpha) \sum_j u_{klj} \leq 0 \quad , \quad v_{kli}, u_{klj} \geq \epsilon x_{kl} \\
 & x_{kl}, y_k \in \{0, 1\}
 \end{aligned}$$

مدل تلفیقی مسئله تخصیص مراکز سرویس با ظرفیت محدود و تحلیل پوششی داده‌ها با در نظر گرفتن وزن نسبی شاخص‌ها به شکل زیر خواهد بود:

$$\begin{aligned}
 & \max \quad \sum_k \sum_l \sum_j u_{klj} O_{jkl} \\
 & \min \quad \sum_k \sum_l c_{kl} b_{kl} + \sum_k F_k y_k \\
 & s.t. \\
 & b_{kl} \leq \min[dem_l, O_{mkl}] y_k, \quad \sum_k b_{kl} = dem_l \\
 & u_{klj} O_{jkl} \leq x_{kl} \leq b_{kl}, \quad x_{kl} \leq y_k \\
 & \sum_j u_{klj} O_{jrs} - \sum_i v_{kli} I_{irs} \leq 0, \quad \sum_k x_{kl} \geq 1 \\
 & (w_{kli} \alpha \sum_i v_{kli}) - v_{kli} \leq 0, \quad \sum_i v_{kli} I_{ikl} = x_{kl} \\
 & v_{kli} - x_{kl} (1 - \alpha + w_{kli} \alpha) \sum_i v_{kli} \leq 0, \quad (w_{klj} \alpha \sum_j u_{klj}) - u_{klj} \leq 0 \\
 & u_{klj} - x_{kl} (1 - \alpha + w_{klj} \alpha) \sum_j u_{klj} \leq 0, \quad v_{kli}, u_{klj} \geq \epsilon x_{kl} \\
 & b_{kl} \geq 0, \quad x_{kl}, y_k \in \{0, 1\}
 \end{aligned} \tag{۸-۴}$$

با دقت در دو مدل ارائه شده و محدودیت‌های اضافه شده به این دو مدل در می‌یابیم که متغیر تصمیم x_{kl} در محدودیت‌های جدید به گونه‌ای اعمال شده است که در صورتی که سرویس دهنده‌ی k به سرویس گیرنده l سرویس ندهد (یعنی $x_{kl} = 0$) وزن شاخص‌های واحد kl (سرویس دهنده‌ی k - سرویس گیرنده‌ی l) توسط پردازشگر صفر در نظر گرفته می‌شود. در حقیقت همان رویکردی که در تلفیق اولیه مد نظر بود (صفر شدن وزن شاخص‌های واحدهایی که عمل سرویس‌دهی را انجام نمی‌دهد) در افزودن این محدودیت‌های کنترل وزن نسبی به مدل رعایت شده است. همان طور که قبلاً گفته شد α درجه‌ی تطابق وزن نسبی شاخص‌ها حاصل از حل مدل با نظر کارشناسان است که مقداری از بازه‌ی بسته صفر و یک می‌پذیرد و باید قبل از حل مدل توسط تصمیم‌گیرنده مشخص شود. با اضافه شدن محدودیت‌های کنترل وزنی به مدل در حقیقت نظر کارشناسان، در انتخاب مجموعه سرویس دهنده‌های بهین اعمال می‌شود و به عبارتی کیفیت انتخاب بالا می‌رود. حل مدل تلفیقی ارائه شده در این بخش به ما مجموعه‌ای از سرویس دهنده‌ها را خواهد داد که علاوه بر کمترین هزینه بالاترین میزان کارایی را با لحاظ شدن نظر

کارشناسان در مورد وزن شاخص‌ها (به کمک محدودیت‌های کنترل وزن نسبی) دارا هستند.

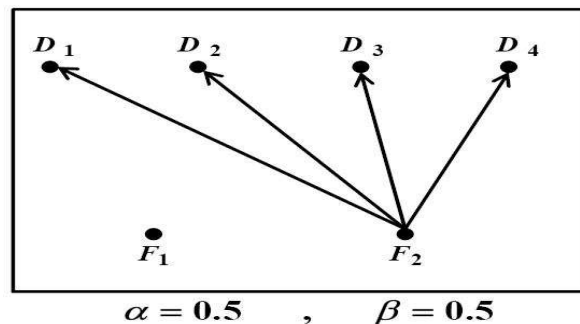
۴-۵ یک مدل پیشنهادی برای مسئله‌ی مکانیابی-تحلیل پوششی داده‌ها

مدلی که در بخش‌های قبل برای مسئله تخصیص مراکز سرویس با ظرفیت سرویس نامحدود و تحلیل پوششی داده‌ها ارائه شده بود به دلیل اختلاف زیاد بین جواب بهینه تابع هدف اول (ماکسیم‌سازی کارایی) و تابع هدف دوم (مینیم‌سازی هزینه) نواقصی بر آن وارد است؛ در حقیقت این اختلاف زیاد مانع از آن می‌گردد که میزان اهمیت هر یک از توابع هدف در جواب نهایی به وضوح مشاهده گردد. ما در مدلی که در این بخش ارائه می‌دهیم این نقص را برطرف می‌سازیم. در ابتدا با توجه به شاخص‌هایی که برای هر واحد (سرویس‌دهنده-سرویس‌گیرنده) تعریف کرده‌ایم میزان کارایی آن واحد را بدون در نظر گرفتن هزینه‌ی سرویس محاسبه و آن را با E نشان می‌دهیم (E_{kl} کارایی سرویس‌دهنده‌ی k وقتی به سرویس گیرنده‌ی l سرویس می‌دهد) در ادامه میزان هزینه را برای هر جفت سرویس‌دهنده-سرویس‌گیرنده بدون در نظر گرفتن کارایی واحد محاسبه می‌کنیم و آن را با W نشان می‌دهیم (W_{kl} هزینه‌ی کل سرویسی که سرویس‌دهنده k به سرویس‌گیرنده‌ی l می‌دهد). در حقیقت برای آنکه مقدار کارایی و هزینه هر واحد هم پایه باشند (دارای یک درجه‌ی بزرگی باشند) مقدار کارایی و هزینه هر واحد را بعد از تلفیق این دو مسئله بر مقدار E_{kl} و W_{kl} مربوط به آن واحد تقسیم می‌کنیم. ما در این مدل اهمیت برآورده شدن هر یک از اهداف (ماکسیم‌سازی کارایی و مینیم‌سازی هزینه) را با قرار دادن ضریب برای متغیرهای تابع هدف نشان دادیم. این ضرائب قبل از حل مدل باید مشخص شوند. مدل تلفیقی تخصیص مراکز

سرویس با ظرفیت سرویس نامحدود و تحلیل پوششی داده‌ها به شکل زیر است :

$$\begin{aligned}
 & \max \quad (\alpha\lambda - \beta\mu) \\
 & s.t \\
 & \lambda = \sum_k \sum_l \frac{\sum_j u_{klj} O_{jkl}}{E_{kl}}, \quad \mu = \sum_k \sum_l \frac{c_{kl} \text{dem}_l x_{kl} + F_k y_k}{w_{kl}} \\
 & u_{klj} O_{jkl} \leq x_{kl}, \quad x_{kl} \leq y_k \\
 & \sum_j u_{klj} O_{jrs} - \sum_i v_{kli} I_{irs} \leq 0, \quad \sum_k x_{kl} = 1 \\
 & v_{kli}, u_{klj} \geq \epsilon x_{kl}, \quad \sum_i v_{kli} I_{ikl} = x_{kl} \\
 & u_{klj}, v_{kli}, \lambda, \mu \geq 0, \quad x_{kl}, y_k \in \{0, 1\}
 \end{aligned} \tag{۹-۴}$$

پارامترها و متغیرهای به کار گرفته شده در مدل (۹-۴) همان پارامترها و متغیرهای مدل اصلی هستند که در بخش‌های قبل معرفی شده‌اند، محدودیت‌ها هم بجزء دو محدودیت اول (که در بالا توضیح داده شد) بقیه همان محدودیت‌های مدل اصلی هستند که از توضیح دوباره آنها صرف نظر می‌کنیم. ضرایب α و β همان ضرایب تعیین اهمیت هریک از اهداف هستند که باید طوری در نظر گرفته شوند که مجموع آنها یک گردد ($\alpha + \beta = 1, \alpha, \beta \geq 0$). متغیر μ از آنجا که از جنس هزینه است و ما سعی در کمینه کردن آن داریم و تابع هدف، ماکسیم‌سازی است با علامت منفی در تابع هدف ظاهر شده‌است. مدلی که تا قبل از این برای مسئله تخصیص مراکز سرویس با ظرفیت سرویس نامحدود و تحلیل پوششی داده‌ها مطرح شده بود یک مدل دو سطحی بود که در آن تخصیص به شکلی صورت می‌گرفت که در سطح اول کارایی و در سطح دوم هزینه بهینه گردد حال ما با وارد کردن این دو هدف در محدودیت‌ها مدل را به یک مدل یک سطحی تبدیل کرده‌ایم. در مدل پیشین مقدار بهینه دو تابع هدف در یک بازه قرار نداشته‌اند و هم بعد نبودند و اختلاف فاحشی که عملاً در مقادیر رخ می‌داد باعث می‌شد جواب خوبی را که منطبق با واقعیت باشد، حاصل نگردد که این نقض در مدل پیشنهادی برطرف گردیده است.



شکل ۴-۴: نمایش گرافیکی تخصیص با رویکرد هزینه و کارایی

۴-۵-۱ مثال عددی

سیستمی را در نظر می‌گیریم که شامل دو سرویس‌دهنده و چهار سرویس‌گیرنده می‌باشد. از آنجا که هر زوج (سرویس‌دهنده-سرویس‌گیرنده) را به عنوان یک واحد تصمیم‌گیری در نظر گرفته‌ایم پس هشت واحد تصمیم‌گیری خواهیم داشت. برای هر واحد سه ورودی و دو خروجی و $\epsilon = 10^{-4}$ در نظر می‌گیریم. مقدار پارامترهای مربوط به ورودی و خروجی واحدها و هزینه‌ی هر واحد سرویس در جدول (۱.۴)، میزان سرویس مورد نیاز برای هر سرویس‌گیرنده و هزینه‌ی استقرار مراکز برای هر سرویس‌دهنده در جدول (۲.۴) آمده است. مقدار کارایی هر واحد بدون در نظر گرفتن بحث هزینه و مقدار هزینه‌ی کل برای هر واحد بدون در نظر گرفتن بحث کارایی با توجه به داده‌های ارائه‌شده محاسبه شده و در جدول (۱۱.۴) آمده است (F_k معرف سرویس‌دهنده‌ی k -ام و D_l معرف سرویس‌گیرنده‌ی l -ام است). نتایج حاصل از حل مدل تلفیقی به روش وزن‌دهی به اهداف با در نظر گرفتن وزن $\alpha, \beta = 0.5$ به شرح جدول (۱۲.۴) می‌باشد. نمایش گرافیکی تخصیص با رویکرد هزینه و کارایی در شکل (۴-۴) نمایش داده شده است.

جدول ۱۱.۴: محاسبه‌ی هزینه و کارایی به شکل مستقل

هزینه	کارایی	سرویس گیرنده	سرویس دهنده
۶۱۸	۰/۸۲	D_1	F_1
۲۱۸۲	۱	D_2	F_1
۱۴۹۰	۰/۹۴	D_3	F_1
۹۱۳	۰/۸۷	D_4	F_1
۴۵۸	۰/۷۲	D_1	F_2
۲۳۸۶	۱	D_2	F_2
۱۸۵۰	۱	D_3	F_2
۱۸۴۹	۱	D_4	F_2

جدول ۱۲.۴: نتایج حل مثال (۴-۵-۱) با در نظر گرفتن $\alpha, \beta = ۰.۵$

تابع هدف کل	تابع هدف کارایی	تابع هدف هزینه	y_2	y_1	x_{24}	x_{23}	x_{22}	x_{21}
-۰/۶	۴	۵/۲۱	۱	۰	۱	۱	۱	۱

فصل ۵

نتیجه‌گیری، پیشنهادات و کارهای آینده

در سال‌های اخیر انواع زیادی از مسائل مکانیابی برای یافتن مجموعه‌ی بهینه‌ی مراکز سرویس برای سرویس‌دهنده‌ها مطرح گردیده است. اخیراً تلفیق این مسائل با روش تحلیل پوششی داده‌ها مورد توجه قرار گرفته است؛ بدین صورت که علاوه بر پیدا کردن مکان بهینه‌ی مراکز سرویس، ما بدنبال کاراترین مراکز نیز می‌باشیم. ما معتقدیم مطرح کردن بحث کارایی، همزمان با مینیم کردن هزینه، مدل حاصل را غنی‌تر و کاراتر می‌سازد. مدل تلفیقی که توسط کلیمبرگ و همکارانش [۲۲] ارائه شده بود در حالت‌هایی که شاخص‌ها بصورت بازه‌ای توسط تصمیم‌گیرنده مطرح شود و یا زمانی که تصمیم‌گیرنده خواهان دخیل شدن نظر کارشناسان در وزن‌های حاصل از حل مدل باشد و یا زمانی که تصادفی بودن شاخص‌ها و در نظر گرفتن خطای تصادفی در اندازه‌گیری کارایی مطرح باشد، دیگر کارآمد نمی‌باشد. ما در این تحقیق مدل تخصیص مراکز سرویس با ظرفیت سرویس محدود—نامحدود را با حالات مختلف تحلیل پوششی داده‌ها تلفیق کردیم تا در مکانیابی مراکز سرویس در شرایط مختلف، کارایی آنها را محاسبه کرده و جواب منطبق‌تری با واقعیت بدست آوریم.

در تمام مثال‌هایی که با دو رویکرد کارایی و هزینه حل شده است همانطور که انتظار می‌رفت وقتی که مثال را با رویکرد هزینه حل می‌کنیم مقدار بهینه‌ی هزینه کمتر از مقدار بهینه‌ی هزینه در زمانی است که مثال را با رویکرد کارایی حل می‌کنیم. این مطلب برای مقدار بهینه‌ی کارایی هم صادق است. نتایج حاصل از حل مثال در حالت تصادفی بسیار نزدیک به نتیجه‌ی حاصل از حل مثال در حالت قطعی است که این ناشی از در نظر گرفتن مقدار ۵٪ برای انحراف مورد قبول از کارایی مورد انتظار می‌باشد. بدون شک هر چه این مقدار بزرگتر باشد تفاوت بین نتایج بیشتر می‌شود.

در حالتی که مثال‌ها با رویکرد کارایی حل شده است تخصیص به صورتی انجام گرفته است که واحدهای منتخب دارای کارایی یک باشند و در حالت هزینه نیز به همین شکل می‌باشد یعنی تخصیص به شکلی صورت گرفته است که کمترین هزینه را داشته باشیم. به دلیل محدودیت‌های موجود در استفاده از

نرم افزارهای تحقیق در عملیات در تعداد محدودیت‌ها و متغیرها، مجبور به طرح مثالی در مقیاس کوچک شدیم. با توجه به اینکه دخیل کردن کارایی واحدها به کمک تحلیل پوششی داده‌ها در مبحث مکانیابی کاری نو و تازه است، می‌توان این مدل را روی بسیاری از مسائل که تاکنون با این نگاه به آنها پرداخته نشده است پیاده‌سازی کرد. همانطور که در این پایان نامه آورده شده است تاکنون فقط روی دو مدل مکانیابی (UFLP-CFLP) و تلفیق آنها با مدل (CCR) کار شده است. نویسنده‌ی پایان نامه بر آن است که در آینده این ایده را برای دیگر مسائل مکانیابی همچون p -میانه و یا p -مرکز و مسئله‌ی پوشش پیاده‌سازی کرده و یا مدل تحلیل پوششی داده‌های بکاررفته را تغییر داده و از مدل‌هایی با «بازده به مقیاس تولید متغیر» همچون (BCC) برای تلفیق استفاده کند. به نظر می‌رسد مدل‌های ارائه شده در این پایان نامه بقدری پویا، کارا و واقعی هستند که در صورت فراهم شدن شرایط می‌توان برای مطالعه‌ی موردی روی یک مسئله‌ی واقعی برای بازنگری در بقا یا احداث واحدهای تجاری، صنعتی و یا خدماتی از آنها استفاده کرد. البته این موضوع نیاز به شناخت و فهم درست مسئله و شرایط مسئله دارد.

پیوست A

مراجع

کتاب نامه

- [1] Banker, R.D., A. Charnes and W.W. Cooper, (1984), "Some models for estimating technical and scale inefficiencies in Data Envelopment Analysis", *Management Science*, Vol. 30, pp: 1261- 1264.
- [2] Bowlin, W.F., A. Charnes, W.W. Cooper, H.D. Sherman, (1985), "Data envelopment analysis and regression approaches to efficiency estimation and evaluation" *Annal Operation Research*, Vol.2, pp: 113-138.
- [3] Charnes, A., W.W. Cooper and E. Rhodes, (1978), "Measuring the efficiency of decision making units," *European Journal of Operational Research*, Vol. 2, pp:429-444.
- [4] Charnes, A., W.W. Cooper and E. Rhodes, (1963), "Deterministic equivalents for optimizing and satisficing under chance constraints" *Operations Research*, Vol.11, 1, pp: 18-39.
- [5] Cooper, W.W., Deng, H., Huang, Z.M., Li, Sx (2004). "Chance Constrained programming approach to congestion in stochastic data envelopment analysis" *European Journal of Operational Research*, Vol. 155,issue pp: 487-501.

-
- [6] Cooper, W.W., Huang, Z.M., Li, Sx (2002). "Chance constrained programming approach to technical efficiencies", *Journal of the Operational Research Society*, No 53, pp: 1347-1354.
 - [7] Cooper, W.W., Huang, Z.M., Li, Sx (1996). "Satisficing dEA models under chance constraints", *Annals of Operation Research*, No 66, issue pp: 279-296.
 - [8] Cooper, W.W., Seiford, L.M., Tone, K., (2007), "*Data Envelopment Analysis*", Second Edition, Springer.
 - [9] Despotis, D.K., Smirlis, Y.G. (2002), "Data envelopment analysis with inexact data", *European Journal of Operational Research* vol. 140, pp: 24-36.
 - [10] Drezner Z., H. Hamacher,. (1995) *Facility location: A Survey of Application and Methods*. Springer-Verlage, Berlin.
 - [11] Dyson, R. G. and E. Thanassoulis, (1988), "Reducing weight flexibility in data envelopment analysis," *Journal of Operational Research Society*, Vol 39, pp: 563-576.
 - [12] Farrell M.J., (1957) "The measurement of productive efficiency", *Journal of Royal Statistical Society* 12, pp: 253-281.
 - [13] Fethi, M.D., Jackson, P.M., Weyman-Jones, T.G. (2001). "An Empirical study stochastic DEA and financial performance: the case of the Turkish commercial banking industry", *INFORMS International Hawaii Conference*, Maui, Hawaii, USA, June 17-20.
 - [14] Francis, R. L., L. F. McGinnis, Jr., and J. A. White. (1992) "*Facility Layout and Location: An Analytical Approach, 2nd ed.*," Prentice Hall.

-
- [15] Handler, G.Y., Mirchandani, P.B., (1979). "*Location on Networks: Theory and Algorithms*." M.I.T. Press, Cambridge, MA.
- [16] Hakimi, S.L., (1964). "Optimal location of switching center and the absolute centers and the absolute medians of a graph". *Operations Research* 12, 450-459.
- [17] Hakimi, S.L., (1965). "Optimal distribution of switching centers in a communication network and some related theoretic graph theoretic problems". *Operation Research* 13, 462-475.
- [18] Hamacher H. W. and Nickel S. (1998) "Classification of location models" *Location Science*, 6, pp 229-242.
- [19] Hansen, P., M. Labbe, D. Peeters, and J. F. Thisse. (1987) "Single facility location on networks." *Annals of Discrete Math* 31, 113-146.
- [20] Hossainzadeh Lotfi, F. Jahanshahloo, G.R. Rezaei Balf, F, (2007). "Ranking of DMUs on interval data by DEA". *Int. J. Contemp. Math. Science*. Vol 20, pp: 159-166.
- [21] Kanstantinos P. Traintis, Barbara J. Koopes, C. Patrick Koelling, "*Modeling undesirable Outputs in Data Envelopment Analysis: Various Approaches*", 2002
- [22] Klimberg, R.K., Ratick, S.J. (2008), "Modeling data envelopment analysis (DEA) efficient location/allocation decisions", *Computers and Operations Research*, 35, pp: 457-474.
- [23] Klose, A., Drexl, A., (2005), "Facility location models for distribution system design", *European Journal of Operational Research* Vol 162, pp: 4-29.

-
- [24] Krarup, J. and P. M. Pruzan, (1983) "The simple plant location problem: Survey and synthesis." *European Journal of Operational Research* 12, 36-81.
- [25] Land, K., Lovell, C.A.K., Theor, S. (1993), "Chance constrained data envelopment analysis", *Manegerial in and Decisional Economics*, No 14, pp: 541-544.
- [26] Mirchandani, P. B. and R. Francis. (1990). *Discrete Location Theory*. J.Wiley.
- [27] Olesen, O.B. and Petersen, N.C. (1995), "Chance constrained effeciency evaluation", *Management Science*, 41, pp: 442-457.
- [28] Owen, S. H. and M. S. Daskin, (1998). "Strategic facility location: A review." *European Journal of Operational Research* 111, 423-447.
- [29] Panos M. Pardalos, Mauricio G.C. Resende, (2002). "*Handbook of Applied Optimization*", Oxford University Press.
- [30] Ramanathan, R (2006). "Data envelopment analysis for weight derivation and aggregation in the analytic hierarchy process" *Computer and Operation Research*. Vol. 33, pp: 289-307.
- [31] ReVelle, C.S, Swain, R.W., (1970) . "Central facilities location." *Geographical Analysis* 2, 30-42.
- [32] Roll Y., W. Cook and B. Golany, (1991), "Controlling factor weights in DEA", *IIE Trans.*, Vol. 23, pp: 2-9.
- [33] Roll Y, and B. Golany, (1993), "Alternate method of treating factor weights in DEA", *Omega* Vol. 21, pp: 99-109.

-
- [34] Saaty, T.L. (1980), "The analytical hierarchy process, Pknning, Priority, Resource Allocation", *RWS Publication, USA*.
- [35] Scaparra, M. P., and M. G. Scutella. "Facilities, locations, customers: Building blocks of location models: A survey." *Thechnical Report: tr-0118*, University of Piza, Italy: 2001.
- [36] Stern, Z.S. Mehrez, A. Hadad, Y, (2000), "An AHP/DEA methodology for ranking decision making units". *Intl.Trans. in Op. Res.* Vol. 7, pp: 109-124.
- [37] Sueyoshi, Toshiyuki. (2000). "Stochastic DEA for Restructure Strategy: an Application to a Japanes Petroleum company", *Omega*, No. 28, pp: 385-398.
- [38] Wang, Y.M. Luo, Y. Elhag, T.M, (2008). "A integrated AHP-DEA methodology for bridge risk assessment". *Computer and Industrial Engineering.* VOL. 54, pp: 513-525.
- [39] Wang, Y.M. Parken, C. Luo, Y, (2007). "A linear programming method for generating the most favorable weights. from a pairwise comparison matrix" *Computer Operation Research.* VOL. 35, pp: 3918-3930.
- [40] Weber A., (1926). *Uber den Standort der Industrient.* Tubingen, 1909; English Trans.: *Theory of Location of Industries*, (C.J.Friedrich, ed and trans.), Chicago University Press, Chicago, Illinois.
- [41] Yang, T. Kup, C, (2003). "A hierarchical AHP/DEA methodology for the facilities layout design problem". *European Journal of operational Research.* Vol . 147, pp: 128-136.
- [42] Zadeh, L.A., (1965). "Fuzzy set" ,*Information and Contorol* 8 338-353.

[۴۳] الفت، ل.، آرمان، م. ح.، (۱۳۸۵)، «اعمال محدودیت‌های کنترل وزن نسبی ورودی‌ها و

خروجی‌ها در مدل‌های DEA»، چهارمین کنفرانس بین‌المللی مدیریت.

[۴۴] اصغری‌پور، م. ج.، (۱۳۸۵)، «تصمیم‌گیری چند معیاره»، چاپ چهارم، انتشارات دانشگاه

تهران.

[۴۵] جوادی، ج.، (۱۳۷۶)، «DEA محدود شده به قیود تصادفی»، پایان نامه کارشناسی

ارشد، دانشگاه تربیت معلم تهران.

[۴۶] جهان‌شاهلو، غ.، حسین‌زاده لطفی، ف.، نیکو مرام، ه.، (۱۳۸۷)، «تحلیل پوششی داده‌ها و

کاربردهای آن»، چاپ اول، دانشگاه آزاد اسلامی واحد علوم تحقیقات.

[۴۷] شهرپاری، س.، (۱۳۸۲)، «ارثه‌ی یک مدل DEA فازی برای ارزیابی عملکرد نسبی

دانشکده‌های علوم انسانی دانشگاه تهران»، پایان نامه کارشناسی ارشد، انتشارات دانشگاه

تهران.

[۴۸] صارمی، م.، خونی، ا.، (۱۳۸۳)، «تعیین و پیش‌بینی کارایی شعب بانک ملت استان

قزوین با استفاده از روش تحلیل پوششی داده‌های تصادفی»، دانش مدیریت، شماره ۶۴،

صص: ۱۲۶-۱۰۷.

[۴۹] کوچک‌زاده، ا.، (۱۳۸۴)، «حل مدل BBC فازی در مدل تحلیل پوششی داده‌ها»،

چهارمین کنفرانس بین‌المللی مهندسی صنایع.

[۵۰] فروند، ج.، (۱۳۷۸)، «آمار ریاضی»، عمیدی، ع.، وحیدی‌اصل، م. ق.، چاپ سوم، مرکز

نشر دانشگاهی.

- [۵۱] فرهانی زنجیرانی، ر.، (۱۳۸۵)، «جزوه‌ی درسی طراحی سیستم‌های صنعتی (مکان‌یابی مراکز)»، دانشگاه امیرکبیر.
- [۵۲] فضلی، ص.، آذر، ع.، (۱۳۸۱)، «طراحی مدل ریاضی ارزیابی عملکرد مدیر با استفاده از تحلیل پوششی داده‌ها DEA»، دوره ۶، شماره ۳، دانشگاه تربیت مدرس
- [۵۳] قدسی‌پور، س.ح.، (۱۳۸۱)، «فرایند تحلیل سلسله‌مراتبی (AHP)»، چاپ سوم، مرکز نشر دانشگاهی دانشگاه صنعتی امیرکبیر.
- [۵۴] قصیری، ک.، مهرنو، ح.، جعفریان مقدم، ا.ر.، (۱۳۸۶)، «مقدمه‌ای بر تحلیل پوششی داده‌های فازی»، چاپ اول، مرکز انتشارات علمی دانشگاه آزاد اسلامی واحد قزوین.
- [۵۵] عیوض‌پور، ج.، (۱۳۸۵)، پایان‌نامه کارشناسی ارشد: «مسیریابی وسائل نقلیه با استفاده از الگوریتم فوق‌ابتکاری Tabu Search»، دانشگاه تربیت مدرس.
- [۵۶] معماریانی، ع.، ساعتی‌مهندی، ص.، (۱۳۸۱)، «نظریه مجموعه‌های فازی و تحلیل پوششی داده‌ها»، سومین همایش ریاضی و کاربردهای آن، گروه ریاضی دانشگاه سیستان و بلوچستان، ص ۱۷۰-۱۵۹.
- [۵۷] مهرگان، م.ر.، (۱۳۸۳)، «پژوهش عملیاتی پیشرفته»، چاپ اول، نشر کتاب دانشگاهی، ص ۱۹۱-۱۸۲.
- [۵۸] میرجلالی، ف.، (۱۳۸۷)، «مسئله‌ی مرکز با وزن مثبت و منفی روی شبکه»، پایان‌نامه کارشناسی ارشد، دانشگاه صنعتی شاهرود.

[۵۹] نورا، ع.، (۱۳۸۳)، «روش حلی برای مسائل برنامه‌ریزی خطی بازه‌ای و برنامه‌ریزی

خطی فازی»، پنجمین کنفرانس سیستم‌های فازی ایران.

[۶۰] وینستون، ول.، (۱۳۸۳) «تحقیق در عملیات: برنامه‌ریزی خطی»، میرحسینی، س.ع.،

علیرضایی، م.ر.، انتشارات دانشگاه صنعتی شاهرود.

پیوست B

منطق فازی

تئوری مجموعه‌های فازی و منطق فازی را اولین بار پرفسور لطفی زاده^۱ در مقاله ای بنام «مجموعه‌های فازی، اطلاعات و کنترل» در سال ۱۹۶۵ معرفی نمود. هدف اولیه او در آن زمان، توسعه مدلی کارآمدتر برای توصیف فرایند پردازش زبان‌های طبیعی بود. او مفاهیم و اطلاعاتی همچون مجموعه‌های فازی، رویدادهای فازی، اعداد فازی و فازی سازی را وارد علوم ریاضیات و مهندسی نمود [۴۲].

ریاضیات فازی یک فرا مجموعه از منطق دودوئی است که بر مفهوم درستی نسبی، دلالت می‌کند. منطق کلاسیک هر چیزی را بر اساس یک سیستم دوتایی نشان می‌دهد (درست یا غلط، ۰ یا ۱، سیاه یا سفید) ولی منطق فازی درستی هر چیزی را با یک عدد که مقدار آن بین صفر و یک است نشان می‌دهد. مثلاً اگر رنگ سیاه را عدد صفر و رنگ سفید را با عدد یک نشان دهیم، آنگاه رنگ خاکستری عددی بین صفر و یک خواهد بود. منطق فازی معتقد است که ابهام در ماهیت علم است. بر خلاف دیگران که معتقد بودند که باید تقریب‌ها را دقیق‌تر کرد تا بهره‌وری افزایش یابد، در منطق ارسطویی، یک دسته بندی درست و نادرست وجود دارد. تمام گزاره‌ها درست یا نادرست هستند. برای مثال جمله «هوا سرد است» در منطق ارسطویی اساساً یک گزاره نمی‌باشد، چرا که مقدار سرد بودن برای افراد مختلف متفاوت است و این جمله اساساً همیشه درست یا همیشه نادرست نیست.

در منطق فازی جملاتی هستند که درستی آن گاهی کم و گاهی زیاد است، گاهی همیشه درست و گاهی همیشه نادرست و گاهی تا حدودی درست است. بنیاد منطق فازی بر شالوده نظریه مجموعه‌های فازی استوار است. این نظریه تعمیمی از نظریه کلاسیک مجموعه‌ها در علوم ریاضی است. در تئوری کلاسیک مجموعه‌ها، یک عنصر یا عضو مجموعه هست یا نیست در حقیقت عضویت عناصر از یک الگوی صفر و یک (دودوئی) تبعیت می‌کند. اما تئوری مجموعه‌های فازی این مفهوم را بسط می‌دهد و عضویت درجه بندی شده را مطرح می‌کند. به این ترتیب یک عنصر می‌تواند تا درجاتی و نه کاملاً عضو یک

مجموعه باشد.

تعریف ۱.B اگر X یک مجموعه باشد آنگاه مجموعه فازی $\bar{A}$ در X را به صورت زیر نمایش می دهیم.

$$\bar{A} = \{(x, \mu_{\bar{A}}(x)) \mid x \in X\}$$

که $\mu_{\bar{A}}(x)$ را تابع عضویت یا تابع سازگاری می نامند و به صورت زیر تعریف می شود.

$$\mu_{\bar{A}}(x) : X \rightarrow [0, 1]$$

که صفر به این معنا است که شیء مورد نظر عضو مجموعه نیست و یک به معنای عضویت کامل شیء نسبت به مجموعه است.

تعریف ۲.B مجموعه همه عناصری که در میان مجموعه $\bar{A}$ درجه عضویت آنها حد اقل α باشد را مجموعه α -برش ضعیف می نامیم و بصورت زیر نمایش می دهیم:

$$A_{\alpha} = \{x \in X \mid \mu_{\bar{A}}(x) \geq \alpha\}$$

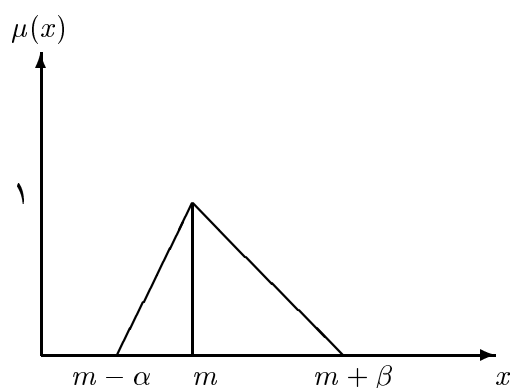
و مجموعه $\bar{A}_{\alpha} = \{x \in X \mid \mu_{\bar{A}}(x) > \alpha\}$ را α -برش قوی می نامیم .

تعریف ۳.B اگر مجموعه فازی A بر روی مجموعه مرجع R از اعداد حقیقی، در شرایط زیر صدق نماید آن را عدد فازی می نامیم.

الف) A یک مجموعه فازی محدب باشد.

ب) یک x_0 موجود باشد که به ازای آن $\mu_A(x_0) = 1$

ج) μ_A پیوسته باشد.



شکل B-۱: تابع عضویت $\tilde{A} = \langle m, \alpha, \beta \rangle$

تعریف ۴.B اگر عدد فازی A شرط زیر را دارا باشد، آن را عدد فازی مسطح می‌نامیم

$$(m_1, m_2) \in R, m_1 \leq m_2; \quad \forall x \in [m_1, m_2] \quad \mu_A(x) = 1.$$

بطور کلی یک عدد فازی مثلثی را با سه‌تایی $\tilde{A} = \langle m, \alpha, \beta \rangle$ و با درجه عضویت زیر نمایش می‌دهیم.

$$\mu(x) = \begin{cases} 0 & \text{for } x \leq m - \alpha \\ 1 - \frac{m-x}{\alpha} & \text{for } m - \alpha < x < m \\ 1 & \text{for } x = m \\ 1 - \frac{x-m}{\beta} & \text{for } m < x < m + \beta \\ 0 & \text{for } x \geq m + \beta \end{cases} \quad (\text{B-1})$$

اعداد فازی مثلثی را با $\tilde{A} = (\underline{a}, a, \bar{a})$ نیز نمایش می‌دهند که $\alpha = a - \underline{a}$ و $\beta = \bar{a} - a$ است. در برخی

تحقیقات، اعداد فازی مثلثی به صورت زیر به بازه‌های مبتنی بر برش α تبدیل شده است.

$$A_\alpha = [\underline{a} + (a - \underline{a})\alpha, \bar{a} - (\bar{a} - a)\alpha]$$

اگر $\tilde{M} = \langle m, \alpha, \beta \rangle$ و $\tilde{N} = \langle n, \gamma, \delta \rangle$ آنگاه عملیات بر روی این دو عدد بصورت زیر است.

مجموع دو عدد $\tilde{M}$ و $\tilde{N}$:

$$\tilde{M} \oplus \tilde{N} = \langle m + n, \alpha + \gamma, \beta + \delta \rangle$$

ضرب دو عدد $\tilde{M}$ و $\tilde{N}$:

$$\tilde{M} \odot \tilde{N} = \begin{cases} \langle mn, m\gamma + n\alpha, m\delta + n\beta \rangle & \text{when } m \geq 0, n \geq 0 \\ \langle mn, n\alpha - m\delta, n\beta - m\gamma \rangle & \text{when } m \leq 0, n \geq 0 \\ \langle mn, -m\delta - n\beta, -m\gamma - n\alpha \rangle & \text{when } m \leq 0, n \leq 0 \end{cases}$$

تفریق دو عدد $\tilde{M}$ و $\tilde{N}$:

$$\tilde{M} \ominus \tilde{N} = \langle m - n, \alpha - \delta, \beta - \gamma \rangle$$

تقسیم دو عدد $\tilde{M}$ و $\tilde{N}$:

$$\tilde{M} \oslash \tilde{N} = \left\langle \frac{m}{n}, \frac{m\delta + n\alpha}{n^2}, \frac{m\gamma + n\beta}{n^2} \right\rangle$$

اگر $\tilde{N}$ یک عدد قطعی مانند K باشد آنگاه

$$\tilde{M} \odot \tilde{N} = \tilde{M}.k = \langle m.k, \alpha.k, \beta.k \rangle \quad \text{به ازای هر } \lambda > 0$$

$$\tilde{M} \odot \tilde{N} = \tilde{M}.k = \langle m.k, -\beta.k, -\alpha.k \rangle \quad \text{به ازای هر } \lambda < 0$$

پیوست C

واژه نامه

واژه نامه‌ی فارسی به انگلیسی

<i>Assessment</i>	ارزیابی
<i>p - center</i>	p -مرکز
<i>p - median</i>	p -میان
<i>Possibility</i>	امکان
<i>Increasing Return to Scale(IRS)</i>	بازده به مقیاس افزایشی
<i>Constant Return to Scale(CRS)</i>	بازده به مقیاس ثابت
<i>Decreasing Return to Scale(DRS)</i>	بازده به مقیاس کاهشی
<i>Linear Programming</i>	برنامه‌ریزی خطی
<i>Multiple Objective Linear Programming</i>	برنامه‌ریزی خطی چند هدفه
<i>Productivity</i>	بهره‌وری
<i>Production Function</i>	تابع تولید
<i>Objective function</i>	تابع هدف
<i>Data Envelopment Analysis(DEA)</i>	تحلیل پوششی داده‌ها
<i>Stochastic Data Envelopment Analysis(SDEA)</i>	تحلیل پوششی داده‌های تصادفی
<i>Allocation</i>	تخصیص
<i>Output</i>	خروجی
<i>Measurment error</i>	خطای اندازه‌گیری

<i>Vertex</i>	رأس
<i>Parametric Methods</i>	روش‌های پارامتری
<i>Non – Parametric Methods</i>	روش‌های غیر پارامتری
<i>Consistency</i>	سازگاری
<i>Facility</i>	سرویس دهنده
<i>Index</i>	شاخص
<i>Network</i>	شبکه
<i>Random Factors</i>	عوامل تصادفی
<i>Imprecise</i>	غیر دقیق
<i>Crisp</i>	قطعی
<i>Subjective Judgment</i>	قضاوت ذهنی
<i>Analytic Hierarchy Process(AHP)</i>	فرایند تحلیل سلسله مراتبی
<i>Efficiency</i>	کارایی
<i>Technical Efficiency</i>	کارایی تکنیکی
<i>Strong Efficiency</i>	کارایی قوی
<i>Relatively Efficiency</i>	کارایی نسبی
<i>Lower bound</i>	کران پائین
<i>Upper bound</i>	کران بالا
<i>Minisum</i>	کمترین مجموع
<i>Node</i>	گره

<i>eigen Value</i>	ماتریس مقایسات
<i>Convex</i>	محدب
<i>Constraint</i>	محدودیت
<i>Weight Restrictions</i>	محدودیت‌های وزنی
<i>Efficient Frontier</i>	مرز کارا
<i>Concave</i>	مقعر
<i>Satisfactory Model</i>	مدل رضایت‌بخشی
<i>Two – Level Mathematic Model</i>	مدل ریاضی دو سطحی
<i>Future Model</i>	مدل آینده‌نگر
<i>Weber problem</i>	مسئله‌ی وبر
<i>Criteria</i>	معیار
<i>Pair wise</i>	مقایسه‌ی زوجی
<i>Location</i>	مکانیابی
<i>Fuzzy logic</i>	منطق فازی
<i>Case study</i>	مطالعه‌ی موردی
<i>Minimax</i>	مینیماکس
<i>Normalization</i>	نرمالیزه کردن
<i>Demand Point</i>	نقطه‌ی تقاضا
<i>Decision Making Unit(DMU)</i>	واحد تصمیم گیرنده
<i>Input</i>	ورودی

Edge پال

پیوست D

فهرست الفبائی

فهرست الفبایی

تحلیل سلسله مراتبی، ۴۵، ۵۳-۵۵، ۵۷-۵۹،	α -برش، ۱۰۱
۶۱	α -برش
تخصیص، ۲، ۱۰، ۹، ۱۳، ۱۸، ۲۳، ۴۷، ۴۶،	ضعیف، ۱۰۱
۵۳، ۶۲، ۷۴، ۷۷، ۷۶، ۸۴-۸۱	قوی، ۱۰۱
سرویس	اعداد فازی مثلثی، ۱۰۲
دهنده، ۶، ۱۱-۹، ۷۶، ۷۸، ۸۳، ۸۲	تابع تولید، ۱۴
گیرنده، ۱۱-۹، ۷۸، ۸۳، ۸۲	تحلیل پوششی داده‌ها، ۱۴، ۱۳، ۱۷، ۱۶، ۲۲،
عدد فازی مسطح، ۱۰۲	۲۵، ۳۸-۳۶، ۴۵، ۴۸، ۴۷، ۵۱-
کارایی، ۱۴، ۲۶-۱۶، ۲۸، ۴۱-۳۶، ۴۴، ۴۷،	۵۳، ۶۰-۵۷، ۷۴، ۷۷، ۸۴-۸۲،
۵۱، ۵۳، ۵۸، ۷۰، ۷۷، ۷۹، ۸۳،	۸۸
۸۴، ۸۸	تحلیل پوششی داده‌های
مجموعه فازی، ۱۰۱	آینده‌نگر، ۴۱، ۷۸، ۷۷
مجموعه فازی	بازه‌ای، ۷۶-۷۴
محدب، ۱۰۱	تصادفی، ۳۸، ۷۷
محدودیت وزنی، ۲۳، ۴۵، ۴۸	قطعی، ۳۵، ۷۹

مکانیابی، ۶-۲، ۱۰-۸، ۵۲، ۶۱، ۸۸

وزن نسبی، ۴۸، ۵۱، ۵۰، ۸۲، ۸۱

هدف های کششی، ۵

هدف های فشاری، ۶

هدف های متعادل، ۶

p-مرکز، ۸

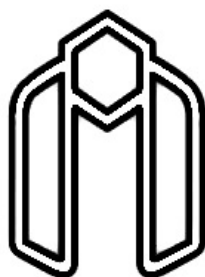
p-میانه، ۸

Abstract

There are many researches on location problem and their extension, also many authors work on data envelopment analysis and its new models. However the hybridation of these two methods which finding facilities with maximum efficiency and minimum cost is new and is introduced by Klimberg and Ratiuk in 2008.

In this thesis we first consider and explain the model of Klimberg and Ratiuk, then present some new models with different conditions. In the thesis we want to assign facilities to clients where the cost of establishing facilities and service is known. We consider each pair facility-client as a decision unit, and define some characteristics to verify efficiency of each unit. The goal of hybridation of location and DEA problems is finding a set of units such that this set has maximum total efficiency and minimum cost service. In this thesis we introduce two models for capacitated and uncapacitated facility location models. We consider the DEA and explain this method in chapter 2. In chapter 3 we introduce an algorithm for finding the weights of vertices as an application of DEA in location theory. This algorithm is a hybridation of DEA and AHP method. In chapter 4 which is the main one we hybrid DEA and location problem and some examples are solved by this method.

Keywords: Location; allocation; DEA; Relatively weight; future DEA model.



Shahrood University of Technology

Faculty of Mathematics

Combination of the location and DEA models

Student: **Shahzad Radkani**

Supervisors:

Dr. Jafar Fathali

Dr. Ahmad Nezakati Reza Zadeh

Jonuary 2010