

ارایه تابع تشخیص q توپین با استفاده از DEA-DA و رگرسیون لجستیک

محمد رحیم رمضانیان^{۱*}، کیخسرو یاکیده^۲، زهرا صفدری سرخزو^۳

۱- دانشیار، دانشگاه گیلان، گروه مدیریت، رشت، ایران

۲- استادیار، دانشگاه گیلان، گروه مدیریت، رشت، ایران

۳- کارشناس ارشد مدیریت صنعتی، دانشگاه گیلان، گروه مدیریت، رشت، ایران

رسید مقاله: ۱۲ مهر ۱۳۹۵

پذیرش مقاله: ۲۰ خرداد ۱۳۹۶

چکیده

تحلیل تشخیص، تشخیص معادله‌ای است مبتنی بر نظم موجود در داده‌ها که براساس آن می‌توان در مورد عضویت یک واحد در یکی از دو یا چند گروه بر مبنای تعدادی متغیر مستقل اظهار نظر کرد. DEA-DA و رگرسیون لجستیک، از جمله روش‌های تحلیل تشخیص هستند که مفروضات پیچیده آماری ندارند و در شرایط بسیاری از پژوهش‌ها قابل کاربردند. از طرفی شاخص q توپین، یکی از شاخص‌های پرکاربرد مالی است که به عنوان معیاری جهت ارزیابی عملکرد به کار می‌رود؛ اما ابهاماتی نیز در دقت محاسبه این شاخص وجود دارد که محققان زیادی در رفع ابهامات این شاخص و بهبود عملکرد آن کوشیده‌اند. پژوهش حاضر نیز به دنبال ارایه تابع تشخیص مناسب شاخص توپین شرکت‌های پذیرفته شده بورس اوراق بهادار تهران است که بر اساس نسبت‌های مالی پرکاربرد مرتبط با این شاخص می‌باشد؛ لذا بعد از یافتن شاخص‌های مرتبط با q توپین، برای داده‌های ۱۸۴ شرکت پذیرفته شده در بورس اوراق بهادار تهران در سال مالی منتهی به ۲۹ اسفند ۱۳۹۳، تحلیل تشخیص پوششی و رگرسیون لجستیک را به طور موازی اجرا کرده و با مقایسه نتایج به دست آمده و حذف تعدادی از شاخص‌ها که در هر دو روش از نقطه نظر ریاضی و آماری وضعیت بی‌ثباتی داشتند، بار دیگر دو روش را با ۵ متغیر مستقل اجرا نموده که وضعیت بهتر و باثبات‌تری در خروجی هر دو روش مشاهده گردید. سپس نتایج گزارش و تحلیل خواهد شد.

کلمات کلیدی: تابع تشخیص، شاخص q توپین، رگرسیون لجستیک، تحلیل تشخیص پوششی.

۱ مقدمه

پیشرفت جوامع و رشد و گسترده‌گی فعالیت‌های مختلف در دوران معاصر موجب با اهمیت‌تر شدن ارزیابی عملکرد شده است. ارزیابی عملکرد یکی از اقدامات اساسی و ضروری، در امر برنامه‌ریزی و هدف‌گذاری مدیران بیان می‌شود، در واقع ارزیابی عملکرد در انتخاب استراتژی و ساختار مالی به مدیران کمک می‌کند و نشان می‌دهد که چگونه استراتژی‌ها و ساختار مالی، بر ارزش بازار سهام اثر می‌گذارد؛ لذا سنجش عملکرد

* عهده‌دار مکاتبات

آدرس الکترونیکی: ramazanian @guilan.ac.ir

شرکت‌ها، ارزیابی کلی از وضعیت مالی و نتایج عملیات برای اخذ تصمیمات منطقی می‌باشد [۱]. استفاده کنندگان گزارش‌های مالی با استفاده از معیارهای مختلف، عملکرد شرکت‌ها را ارزیابی و برآوردی از ارزش شرکت را مشخص می‌کنند [۲]. به عنوان مثال تعدادی از نسبت‌های مالی مانند ROA، ROE وجود دارند که معمولاً با رجوع به آن‌ها می‌توان به وضعیت نسبی شرکت‌ها پی برد؛ اما این نسبت‌ها به دلیل متکی بودن به داده‌های حسابداری که مربوط به دوره‌های گذشته هستند، نمی‌توانند خیلی دقیق باشند.

لذا محققان بازارهای مالی برای ارزیابی اهمیت وقایعی که در شرکت‌ها روی می‌دهد و اثراتی که این وقایع بر شرکت‌ها می‌گذارد، تمایل دارند که به قیمت‌های بازار اتکا کنند. این امر باعث شده است که در تحقیقات خود به جای سنج‌های حسابداری، از سنج‌های ترکیبی مانند q توین^۱ برای اندازه‌گیری عملکرد شرکت‌ها استفاده کنند [۳]؛ زیرا شاخص توین به علت مبتنی بودن بر قیمت بازار شرکت، منعکس کننده قضاوت بازار در مورد وضعیت شرکت نیز می‌باشد.

نظریه توین یکی از نظریه‌های نئوکلاسیکی سرمایه‌گذاری است که در سال ۱۹۶۹ توسط اقتصاددانی به نام جیمز توین^۲ مطرح گردید. صرف‌نظر از اهداف اولیه این نظریه، شاخص توین که نسبت ارزش بازار شرکت به ارزش دفتری یا ارزش جایگزینی دارایی‌های شرکت است، اطلاعات خوبی از وضعیت مالی شرکت یا سهام آن فراهم می‌کند. به عنوان مثال در تحلیل‌های آکادمیک مالی می‌تواند به عنوان یک متغیر کنترل در بررسی اثر نسبت‌های مالی بر یکدیگر، ارزیابی عملکرد مدیران شرکت‌ها [۴]، بررسی تفاوت‌های مقطعی در سرمایه‌گذاری و تنوع تصمیمات [۵] و به عنوان معیاری برای تعیین ارزش شرکت‌ها [۶] استفاده شود.

هرچند شاخص توین از پایه نظری مستحکمی برخوردار است؛ اما محاسبه دقیق آن امکانپذیر نیست؛ زیرا اگر قیمت بازار شرکت مشخص نباشد برخی از ویرایش‌های q توین و اگر قیمت جایگزینی دارایی‌هایی شرکت مشخص نباشد برخی ویرایش‌های دیگر این شاخص قابل محاسبه نیست؛ لذا با توجه به کاربردهای زیاد شاخص توین و از سوی دیگر انتقادات وارد شده به شاخص، منطقی به نظر می‌رسد که برای برآورد شاخص توین از روی دیگر شاخص‌های مالی تلاش شود.

پژوهش‌های زیادی از جمله [۵]، [۷]، [۸] می‌توان یافت که به پیش‌بینی و برآورد این شاخص از روی دیگر نسبت‌های مالی پرداخته‌است؛ اما نکته مشترک همه این پژوهش‌ها این بوده است که با استفاده از روش‌های پیش‌بینی دقیق مانند رگرسیون‌های خطی به پیش‌بینی و برآورد q توین پرداخته‌اند در حالی که شاخص توین یک شاخص پویاست و در لحظه ممکن است مقدار آن دچار تغییر شود؛ لذا این پژوهش به سراغ روش‌های تحلیل تشخیص رفته که به جای تعیین میزان دقیق متغیر وابسته از روی متغیرهای مستقل، به تعیین عضویت گروهی متغیر وابسته می‌پردازد. منظور از تحلیل تشخیص گروه‌بندی داده‌ها به گروه‌های متجانس است به گونه‌ای که مشاهدات هر گروه به یکدیگر بیش‌ترین شباهت را داشته باشد [۹] و چون شاخص توین علی‌رغم پویا بودن،

¹. Tobin's q

². James Tobin

متصور نیست که تغییرات شدیدی داشته باشد؛ لذا منطقی نیست که در مورد یک شاخص پویا از روی دیگر شاخص‌های مالی، قضاوتی محدود در مورد عضویت آن به گروه‌های چندگانه انجام شود.

لذا این پژوهش به کمک اطلاعات به دست آمده از دیگر شاخص‌های مالی ۱۸۴ شرکت پذیرفته شده در بورس اوراق بهادار تهران در سال مالی منتهی به ۲۹ اسفند ۱۳۹۳ و رابطه‌ای که بین این نسبت‌های مالی و شاخص توپین وجود دارد به دنبال بررسی رابطه بین شاخص توپین و تعدادی از دیگر نسبت‌های مالی مرتبط و ارایه تابع تشخیص^۱ مناسب جهت تعیین عضویت گروهی شرکت‌های مورد آزمون برای قضاوت درباره وضعیت شرکت‌ها است.

۲ مبانی نظری و پیشینه پژوهش

۲-۱ شاخص توپین

جیمز توپین برای ارزیابی سودآوری پروژه‌های سرمایه‌گذاری از نسبت زیر به عنوان نسبت q استفاده کرد:

$$q = \frac{\text{در پایان سال (ارزش دفتری جمع بدهی‌های شرکت + ارزش بازار سهام ممتاز + ارزش سهام عادی)}}{\text{ارزش دفتری کل دارایی‌های شرکت در پایان سال}} \quad (*)$$

طبق نظریه توپین، اگر مقدار محاسبه شده برای شاخص q شرکتی بزرگ‌تر از یک باشد، انگیزه زیادی برای سرمایه‌گذاری وجود دارد به عبارتی نسبت q بالا، معمولاً نشان‌دهنده ارزشمندی فرصت‌های سرمایه‌گذاری و رشد شرکت است و اگر نسبت q کوچک‌تر از یک باشد، سرمایه‌گذاری متوقف خواهد شد.

با گذشت زمان، پژوهش‌های زیادی به بررسی دقت محاسبه این شاخص پرداخته و آن را مورد نقد قرار داده است. از جمله انتقادات وارد شده این بود که در مخرج کسر این نسبت، از ارزش دفتری دارایی‌ها استفاده می‌شود که مبتنی بر ارزش‌های تاریخی است و ممکن است تفاوت فاحشی با ارزش روز دارایی‌های شرکت داشته باشد و این موضوع به این دلیل حایز اهمیت است که شاخص توپین در واقع یک شاخص پویاست و در لحظه قابل محاسبه می‌باشد؛ اما وجود مقادیر مبتنی بر اطلاعات گذشته می‌تواند در دقت مقدار آن شبهه ایجاد کند.

بعد از این در سال ۱۹۷۷، جیمز توپین با بررسی نقدهای وارد شده بر شاخص توپین ساده، با ایجاد تعدیلاتی، الگوی جدیدی از شاخص q ، به نام q توپین استاندارد را معرفی کرد که از تقسیم مقدار مندرج در صورت کسر رابطه (*) با ارزش جایگزینی دارایی‌های آن به دست می‌آید. ارزش جایگزینی در واقع ارزش روز کل دارایی‌های شرکت از نظر کارشناسان مربوطه می‌باشد.

بعد از این نیز پژوهش‌گران دیگری، از جمله لیندنبرگ و راس [۱۰]، چانگ و پروت [۱۱] و لی وای لن و بادرنت [۱۲] نیز با انجام تعدیلاتی در شاخص q ، الگوهای دیگری از این شاخص را معرفی کردند که تا حدودی ایرادات الگوی قبلی را برطرف می‌نمود.

¹. Discriminant Function

استیون و کنیث [۱۳] به بحث و بررسی درباره q به عنوان معیاری برای اندازه گیری عملکرد شرکت ها پرداختند و نسخه های مختلف q را مورد بررسی قرار دادند و نتایج پژوهش آن ها نشان داد که میانه و میانگین و انحراف معیار برآوردهای نسخه های متداول q تا حدودی با هم برابر هستند.

حیدرپور [۸] با استفاده از نمونه ای با حجم ۱۰۰ تایی طی سال های (۲۰۰۵-۲۰۱۰) به برآورد فاکتورهای موثر بر شاخص توپین در بورس اوراق بهادار تهران می پردازد و با استفاده از رگرسیون خطی ضریب اهمیت هر یک از این شاخص ها را تعیین می کند.

نباوند [۷] در پژوهش خود به بررسی ارتباط بین q توپین به عنوان معیاری برای ارزیابی عملکرد مبتنی بر اطلاعات بازار و نسبت های مالی متداول P/E ، P/B ^۱، EPS ^۲ و ROA و ROE به عنوان معیارهای ارزیابی عملکرد مبتنی بر اطلاعات تاریخی و حسابداری می پردازد.

۲-۲ عوامل مؤثر یا مرتبط با شاخص توپین (متغیرهای مستقل پژوهش)

در این پژوهش، با مرور ادبیات مربوط به شاخص q توپین، از شاخص های مالی زیر به عنوان عوامل مرتبط و موثر شاخص توپین استفاده می گردد.

لازم به ذکر است که رابطه بین شاخص q توپین با شاخص های $NITR$ ، DPR ، $Gror_1$ در منبع [۱۴] و با شاخص $LISE$ در منبع [۱۵] و با شاخص $TDTA$ در منبع [۱۶] و نیز با شاخص های $CATA$ ، $CLTA$ و $Size$ در منبع [۱۷] بررسی و تأیید شده است. در مورد ارتباط شاخص q توپین با شاخص $WCTA$ پژوهشی وجود نداشت؛ اما از آنجا که این شاخص مانند شاخص $CLTA$ و $CATA$ به نوعی نشان دهنده ساختار سرمایه بوده و در بسیاری از پژوهش های مالی از جمله پژوهش های مربوط به ورشکستگی مانند منبع [۱۶] از آن استفاده شده است، در این پژوهش نیز تاثیر این شاخص در تحلیل تشخیص q توپین بررسی خواهد شد.

نسبت سود خالص به فروش کل ($NITR^3$): این نسبت که از تقسیم سود پس از کسر مالیات به کل فروش محاسبه می شود، سودآوری هر ریال از فروش را نشان می دهد.

نسبت پرداختی سود نقدی (DPR^4): این نسبت نشان دهنده درصدی از سود سالیانه است که به طور نقدی بین سهامداران عادی تقسیم می شود. این نسبت که از تقسیم سود تقسیمی هر سهم به سود هر سهم به دست می آید، هم برای سهامداران و هم برای طلبکاران بسیار مهم است [۱۸].

نرخ رشد فروش یک ساله ($Gror_1^5$): این شاخص جزء معیارهای خوشه رشد بوده و نشان دهنده درجه تغییر در عایدات شرکت در طی دوره زمانی یکساله می باشد.

1. Price/Book Value

2. Earnings Per Share

3. Ratio of net income to total revenue

4. Dividend pay-out ratio

5. Growth rate of total revenue on 1 year

نسبت بدهی‌های بهره‌دار به حقوق صاحبان سهام ($LISE^1$): این نسبت از تقسیم بدهی‌های بهره‌دار شرکت که شامل مجموع تسهیلات مالی بلندمدت شرکت‌ها می‌باشد به جمع حقوق صاحبان سهام در پایان سال مالی به دست می‌آید.

نسبت کل بدهی‌ها به کل دارایی‌ها ($TDTA^2$): این نسبت یکی از نسبت‌های اهرمی مهم است که نشان می‌دهد میزان تامین مالی شرکت از طریق بدهی‌ها چقدر است.

نسبت سرمایه در گردش به کل دارایی‌ها ($WCTA^3$): این نسبت که برابر با حاصل تقسیم تفاضل بین دارایی‌های جاری و بدهی‌های جاری به کل دارایی‌ها شرکت می‌باشد، در دسترس بودن سرمایه در گردش را نشان می‌دهد.

نسبت دارایی‌های جاری به کل دارایی‌ها ($CATA^4$): این نسبت که از تقسیم دارایی‌های جاری به کل دارایی‌ها به دست می‌آید می‌تواند به عنوان معیار اندازه‌گیری راهبرد سرمایه‌گذاری شرکت استفاده شود؛ زیرا افزایش در سطح سرمایه‌گذاری در دارایی‌های جاری نشان می‌دهد که مدیران در مدیریت دارایی‌های جاری، راهبرد محافظه کارانه تر را دنبال می‌کنند.

نسبت بدهی‌های جاری به کل دارایی‌ها ($CLTA^5$): این نسبت از تقسیم بدهی‌های جاری به کل دارایی‌ها به دست می‌آید.

اندازه شرکت (Size): بدیهی است که اندازه شرکت، تعیین‌کننده حجم و گستردگی فعالیت‌های یک شرکت است [۱۹]. در این پژوهش، برای محاسبه اندازه شرکت، از لگاریتم طبیعی مجموع دارایی‌های شرکت در پایان سال مالی استفاده شده است.

۲-۳ تحلیل تشخیص رگرسیونی (رگرسیون لجستیک)

همان‌طور که می‌دانیم برای انجام تحلیل رگرسیون خطی، متغیر وابسته باید کمی و در سطح سنجش فاصله‌ای باشد؛ اما گاهی اوقات اتفاق می‌افتد که متغیر وابسته تحقیق در مقیاس فاصله‌ای / نسبتی نبوده و مقیاس آن به صورت اسمی (دوسطحی یا چندسطحی) است.

از نظر مفهومی، رگرسیون لجستیک و رگرسیون خطی قابل قیاس‌اند. در واقع هر دو معادله‌هایی برای پیش‌بینی فراهم می‌سازند؛ لذا رگرسیون لجستیک شبیه به رگرسیون خطی است با این تفاوت که نحوه محاسبه ضرایب در این دو روش یکسان نمی‌باشد. بدین معنی که رگرسیون لجستیک، به جای حداقل کردن مجذور انحرافات، احتمالی را که یک واقعه رخ می‌دهد حداکثر می‌کند؛ زیرا به دلیل دو سطحی بودن متغیر وابسته در رگرسیون لجستیک از فرض برابری واریانس که زیر بنای رگرسیون خطی است صرف نظر می‌شود [۹].

1. Ratio of liability with interest to shareholder's equity
2. Ratio of total debt to total assets
3. Ratio of working capital to total assets
4. Ratio of current assets to total assets
5. Ratio of current liabilities to total assets

نکته اینجاست که رگرسیون خطی، شرایطی را فراهم می‌کند که در آن متغیرهای مستقل با متغیر وابسته، رابطه ثابتی دارند؛ یعنی مثلاً این مقدار تفاوت در متغیر پیش‌بین، به آن مقدار تفاوت در متغیر وابسته منجر می‌شود. هرگاه یک رابطه ثابت توصیف‌کننده رابطه نباشد، استفاده از حداقل مجذورات، همراه با برازش تابع خطی آن راهبرد قابل قبولی نخواهد داشت و این جایی است که رگرسیون لجستیک، با تابع S شکل آن که متغیرهای پیش‌بین را به احتمال وقوع پیشامدها ربط می‌دهد، از توان پیش‌بینی قابل توجهی برخوردار خواهد بود.

شکل عمومی رگرسیون لجستیک به صورت زیر است:

$$\pi = p(x) = \frac{e^{\alpha + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_k x_k}}{1 + e^{\alpha + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_k x_k}} \quad (1)$$

اما برخلاف رگرسیون خطی، رگرسیون لجستیک مستقیماً متغیر وابسته (Y) را مدل‌سازی نمی‌کند؛ بلکه ابتدا متغیر وابسته را به یک متغیر لجوژیت (لگاریتم طبیعی بخت‌های وقوع یا عدم وقوع متغیر Y) تبدیل می‌کند (یعنی به صورت $Ln(\frac{\pi}{1-\pi})$) و سپس از برآورد حداکثر درست‌نمایی برای برآورد ضرایب استفاده می‌کند.

۲-۴ تحلیل تشخیص پوششی (DEA-DA)

تکنیک DEA تکنیکی ناپارامتریک برای سنجش و ارزیابی کارایی نسبی مجموعه‌ای از پدیده‌ها با ورودی‌ها و خروجی‌های قطعی است. به عبارت دیگر روشی است که برای محاسبه کارایی نسبی واحدهای تصمیم‌گیرنده ایجاد شده است که منابع چندگانه مشابهی را برای تولید خروجی‌های مشابه به کار می‌برند [۲۰].

بحث تحلیل پوششی داده‌ها با تز دکتر «ادوارد روز» تحت راهنمایی «کوپر و چارنز» شروع شد که پیشرفت تحصیلی دانش‌آموزان مدارس آمریکا را در سال ۱۹۷۸ مورد ارزیابی قرار داده بودند. نتایج این مطالعه منجر به چاپ اولین مقاله درباره معرفی عمومی DEA در سال ۱۹۷۸ گردید [۲۱]. آن‌ها در مقاله مذکور که به CCR شهرت یافت، برای تعمیم روش دو ورودی و یک خروجی، فارل از بهینه‌سازی به روش برنامه‌ریزی ریاضی استفاده کردند تا بتوانند کارایی سیستم‌هایی با ورودی‌های چندگانه و خروجی‌های چندگانه^۱ را اندازه‌گیری کنند. مدل CCR در سه فرم کسری، ضربی و پوششی مدل‌سازی شده است. بعد از آن، در سپتامبر ۱۹۸۴ بنکر و چارنز و کوپر مفاهیم و مدل‌های DEA را با مفاهیم جدیدی توسعه دادند که حاصل آن مدل BCC است. این مدل برای اندازه‌گیری و تعیین کارایی واحدها برای بالا بردن اندازه‌ی کارایی و با در نظر گرفتن بازده به مقیاس متغیر مورد استفاده قرار می‌گیرد [۲۲].

بعد از آن، در سال ۱۹۸۵ چارنز و همکارانش «مدل جمعی»^۲ را به عنوان یکی دیگر از مدل‌های اساسی در DEA مطرح کردند. مدل‌های مطرح شده تا قبل از مدل جمعی، که مدل‌های شعاعی نیز نامیده می‌شد، در یکی

1. Multiple Input & Multiple Output
2. Additive Model

از دو فرمت ورودی محور یا خروجی محور ارایه شد؛ اما در مدل‌های جمعی، که جزء مدل‌های غیر شعاعی است، تابع هدف طوری تعریف می‌شود که به‌طور همزمان هم به دنبال حداکثر کردن خروجی‌ها و همچنین حداقل کردن ورودی‌ها می‌باشند. همچنین از دیگر ویژگی‌های مدل‌های جمعی این است که می‌تواند مقادیر منفی نیز در ورودی و یا خروجی‌ها داشته باشند [۲۳] که در بسیاری از پژوهش‌ها و از جمله در زمینه‌های مالی که داده‌های منفی به وفور به چشم می‌خورد، قابل کاربرد است.

لذا با توجه به ویژگی‌های گفته شده، مدل تحلیل تشخیص مبتنی بر تحلیل پوششی داده‌ها نیز بر پایه مدل‌های جمعی تعریف می‌گردد.

حال فرض کنید n واحد تصمیم‌گیرنده هر کدام با k عامل مستقل موجود هستند و Z_{ij} ، فاکتور i ام ($i = 1, 2, \dots, k$) از واحد j ام ($j = 1, 2, \dots, n$) می‌باشد و از تاثیر عامل‌ها یا روابطی که استفاده می‌شود هیچ آگاهی قبلی نداریم، فقط می‌دانیم که عضویت گروه شناخته شده است و معلوم است که n_1 واحد در گروه اول (که با $J \in G_1$ نشان می‌دهیم) و n_2 واحد در گروه دوم (که با $J \in G_2$ نشان می‌دهیم) قرار دارد و $n = n_1 + n_2$ [۲۴].

در اولین مدل $DEA-DA$ ارایه شده توسط سوئیشی [۲۴] طی دو مرحله به طبقه‌بندی واحدهای مورد نظر پرداخته می‌شود.

مرحله اول: حداقل کردن تعداد واحدهای عضو ناحیه همپوشانی^۱

$$\begin{aligned} \text{Min} \quad & \sum_{j \in G_1} S_{1j}^+ + \sum_{j \in G_2} S_{2j}^- \\ \text{s.t.} \quad & \sum_{i=1}^k \alpha_i Z_{ij} + s_{1j}^+ - s_{1j}^- = d, \quad j \in G_1, \\ & \sum_{i=1}^k \alpha_i Z_{ij} + s_{2j}^+ - s_{2j}^- = d - \eta, \quad j \in G_2, \\ & \sum_{i=1}^k \alpha_i = 1, \\ & \text{All slacks} \geq 0, \quad \alpha_i \geq 0, \quad d : \text{unrestricted}. \end{aligned} \quad (2)$$

که در آن، S_{1j}^+ و S_{2j}^- نشان‌دهنده انحرافات مثبت و منفی واحدهای گروه G_1 از تابع تشخیص خطی قطعه‌ای $\sum_{i=1}^k \alpha_i Z_{ij} = d$ می‌باشد و به همین ترتیب، S_{2j}^+ و S_{1j}^- نشان‌دهنده انحرافات مثبت و منفی واحدهای گروه G_2 از مقدار تابع تشخیص خطی قطعه‌ای $\sum_{i=1}^k \alpha_i Z_{ij} = d$ می‌باشد از حل مدل در مرحله اول و به‌دست آوردن α_i^* ، در مرحله دوم، $DEA-DA$ می‌تواند همه واحدهایی را که به ناحیه همپوشانی دو گروه تعلق دارد با استفاده از قاعده زیر به یکی از دو گروه G_1 یا G_2 طبقه‌بندی کند:

$$\begin{aligned} \text{if } \sum_{i=1}^k \alpha_i^* Z_{ij} \geq d^*, j \in G_1 \cap G_r \Rightarrow j \in G_1 \\ \text{if } \sum_{i=1}^k \alpha_i^* Z_{ij} < d^*, j \in G_1 \cap G_r \Rightarrow j \in G_r \end{aligned} \quad (3)$$

نکته قابل ذکر این است که مقدار غیر صفر η برای اجتناب از جواب بدیهی که در آن همه وزن‌ها مساوی با صفر است، به مدل اضافه شده است.

اما در کاربردهای مختلف از این مدل مشخص شد که این مدل توانایی پردازش داده‌های با حجم بالا را ندارد، نمی‌تواند مقادیر منفی بپذیرد و همچنین دو صفحه یا خط تمیز دهنده بین دو گروه در دو مرحله ارایه می‌کند؛ لذا سوئیشی [25] در سال 2001 به ارایه مدل جدیدی از DEA-DA پرداخت که معایب وارد شده بر مدل اولیه را در آن برطرف کرده است.

از طرفی در مدل‌های اولیه DEA-DA ارایه شده توسط سوئیشی که بر مبنای برنامه‌ریزی خطی پایه-گذاری شده بود، هدف حداقل کردن مجموع انحرافات می‌باشد در حالی که هدف توابع تشخیص آماری حداقل کردن تعداد طبقه‌بندی‌های اشتباه می‌باشد نه حداقل کردن انحرافات از مرز تابع تشخیص. از سوی دیگر در مدل‌های مبتنی بر برنامه‌ریزی خطی که با الگوریتم‌هایی مانند سیمپلکس حل می‌شود، مدل نیازمند نقاطی است که مرز بهینه را تشکیل می‌دهد و این موضوع در مرحله طبقه‌بندی داده‌ها می‌تواند مساله‌ساز باشد؛ زیرا گروه‌بندی دقیق نقاطی که روی مرز تفکیک قرار دارد، امکان‌پذیر نخواهد بود و لذا برای طبقه‌بندی نقاط روی مرز به کارگیری روش‌های دیگری الزامی خواهد بود؛ لذا جهت غلبه بر ایرادات وارد شده بر این مدل‌ها و بهتر شدن نتایج طبقه‌بندی، سوئیشی [26] مدل جدیدی از DEA-DA ارایه کرد که مبتنی بر MIP^1 می‌باشد. که میزان طبقه‌بندی درست در آن‌ها بیش‌تر از مدل‌های مبتنی بر برنامه‌ریزی خطی (LP) می‌باشد و همچنین این مدل‌ها نیاز به نقاط مرزی برای تشکیل مرز طبقه‌بندی نخواهند داشت. همچنین باید اذعان کرد در مدل‌های مطرح شده، در تغییر (تبدیل) متغیر λ_i به دو متغیر λ_i^+, λ_i^- سوئیشی فرض کرد که این دو متغیر جدید $(\lambda_i^+, \lambda_i^-)$ هیچگاه به-طور همزمان مثبت نخواهند بود و تعریف ریاضی این فرض (یا شرط) به صورت زیر خواهد بود:

$$\begin{aligned} \text{Min } & \sum_{j \in G_1} S_{1j}^+ + \sum_{j \in G_r} S_{rj}^- \\ \text{s.t. } & \sum_{i=1}^k \lambda_i Z_{ij} + S_{1j}^+ - S_{1j}^- = d + 1, \quad j \in G_1, \\ & \sum_{i=1}^k \lambda_i Z_{ij} + S_{rj}^+ - S_{rj}^- = d, \quad j \in G_r, \\ & \sum_{i=1}^k |\lambda_i| = 1, \\ & \text{All Slacks} \geq 0, \quad \lambda_i, d : \text{unrestricted}. \end{aligned} \quad (4)$$

$$\lambda_i^+ = \frac{(|\lambda_i| + \lambda_i)}{2}, \quad \lambda_i^- = \frac{(|\lambda_i| - \lambda_i)}{2}$$

که در آن :

$$\lambda_i = \lambda_i^+ - \lambda_i^- \quad \text{و} \quad |\lambda_i| = \lambda_i^+ + \lambda_i^-$$

و یک چنین تبدیلی نیاز به شرایط حل غیر خطی $\lambda_i^+ \lambda_i^- = 0$ دارد که تضمین می کند این دو متغیر هیچ وقت به طور همزمان مقدار نخواهند گرفت و حتما یکی از این دو مقدار صفر خواهد بود. که تعریف متغیر باینری در MIP امکان پذیر خواهد بود. این مدل در دو مرحله حل می شود که در ادامه بیان خواهد شد.

مرحله اول

$$\text{Min } S$$

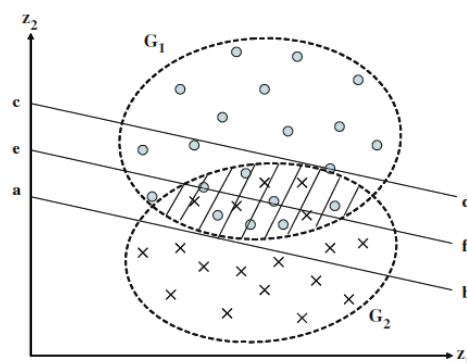
$$\text{s.t.} \quad \sum_{f=1}^h \lambda_f Z_{fj} - d + s \geq 0, \quad j \in G_1,$$

$$\sum_{f=1}^h \lambda_f Z_{fj} - d - s \leq -\varepsilon, \quad j \in G_2, \quad (5)$$

$$\sum_{f=1}^h |\lambda_f| = 1,$$

$$\lambda_f, d, s: \text{URS.}$$

که G_1 و G_2 دو گروه واحدهایی هستند که به دنبال تعیین تعلق واحدها به یکی از این دو گروه هستیم و S اشتراک بین دو گروه، $d+s$ و $d-s$ به ترتیب مرز پایین G_1 و مرز بالای G_2 را مشخص می کنند. λ_f و Z_{fj} به ترتیب شاخص های مالی و ضرایب آن ها در تابع تشخیص را نشان می دهند (شکل ۱).



شکل ۱. گروه بندی واحدها در DEA-DA

اما از آنجا که حل مدل با وجود $\sum_{f=1}^h |\lambda_f| = 1$ مشکل می باشد با انجام تغییراتی مدل به صورت زیر تبدیل می گردد:

$$\begin{aligned}
 & \text{Min } S \\
 & \text{s.t. } \sum_{f=1}^h (\lambda_f^+ - \lambda_f^-) Z_{\bar{f}j} - d + s \geq 0, \quad j \in G_1, \\
 & \quad \sum_{f=1}^h (\lambda_f^+ - \lambda_f^-) Z_{\bar{f}j} - d - s \leq -\varepsilon, \quad j \in G_2, \\
 & \quad \sum_{f=1}^h (\lambda_f^+ - \lambda_f^-) = 1, \\
 & \quad \zeta_f^+ \geq \lambda_f^+ \geq \varepsilon \zeta_f^+ \text{ (for all } f), \quad \zeta_f^- \geq \lambda_f^- \geq \varepsilon \zeta_f^- \text{ (for all } f), \\
 & \quad \zeta_f^+ + \zeta_f^- \leq 1 \text{ (for all } f), \quad \lambda_f^+ + \lambda_f^- \geq \varepsilon \text{ (for all } f), \\
 & \quad \zeta_f^+ \text{ \& } \zeta_f^- : \text{binary, and all other variable } \geq 0. \\
 & \quad d, s : \text{URS}
 \end{aligned} \tag{6}$$

که در آن:

$$\begin{aligned}
 |\lambda_f| &= \lambda_f^+ + \lambda_f^-, \quad (\lambda_f = \lambda_f^+ - \lambda_f^-), \\
 \lambda_f^+ &= (|\lambda_f| + \lambda_f) / 2, \quad \lambda_f^- = (|\lambda_f| - \lambda_f) / 2
 \end{aligned} \tag{7}$$

جواب بهینه این مدل

$$\lambda_f^* = (\lambda_f^{+*} - \lambda_f^{-*}) / 2$$

d^* و s^* می باشد که اگر $s^* \leq 0$ باشد؛ یعنی تمام اعضاء به دو گروه بالا و پایین گروه بندی شده است و داده ای که عضویت آن مشخص نشده باشد وجود ندارد؛ اما $s^* > 0$ نشان دهنده اشتراک بین دو گروه است و برای تعیین تعلق این داده ها باید مدل مرحله دوم را حل کنیم.

مرحله دوم

$$\begin{aligned}
 & \text{Min } \sum_{j \in D_1} y_j + \sum_{j \in D_2} y_j \\
 & \text{s.t. } \sum_{f=1}^h (\lambda_f^+ - \lambda_f^-) z_{\bar{f}j} - c + M y_j \geq 0, \quad j \in D_1, \\
 & \quad \sum_{f=1}^h (\lambda_f^+ - \lambda_f^-) z_{\bar{f}j} - c - M y_j \leq -\varepsilon, \quad j \in D_2, \\
 & \quad \sum_{f=1}^h (\lambda_f^+ - \lambda_f^-) = 1, \\
 & \quad \zeta_f^+ \geq \lambda_f^+ \geq \varepsilon \zeta_f^+ \text{ (for all } f), \quad \zeta_f^- \geq \lambda_f^- \geq \varepsilon \zeta_f^- \text{ (for all } f), \\
 & \quad \zeta_f^+ + \zeta_f^- \leq 1 \text{ (for all } f), \quad \lambda_f^+ + \lambda_f^- \geq \varepsilon \text{ (for all } f), \\
 & \quad c : \text{URS, } \zeta_f^+ \text{ \& } \zeta_f^- : \text{binary, and all other variable } \geq 0.
 \end{aligned} \tag{8}$$

لازم به ذکر است در اینجا، متغیر باینری y_j تعداد واحدهایی از $D_1 \cup D_2$ را که به اشتباه طبقه‌بندی شده‌است، می‌شمارد؛ یعنی تابع هدف مرحله دوم $DEA-DA$ به دنبال حداقل کردن تعداد طبقه‌بندی‌های اشتباه می‌باشد. بعد از به‌دست آوردن λ^* و c^* به‌عنوان جواب بهینه مرحله دوم، طبقه‌بندی داده‌ها به دو گروه G_1 و G_2 بدین-صورت انجام می‌گیرد که اگر

$$\sum_{f=1}^h \lambda_f^* z_{fj} \geq c^* \quad (9)$$

آنگاه j امین مشاهده به G_1 تعلق دارد و اگر

$$\sum_{f=1}^h \lambda_f^* z_{fj} \leq c^* - \varepsilon \quad (10)$$

باشد به G_2 تعلق دارد و بنابراین تمامی مشاهدات به یکی از دو گروه G_1 یا G_2 طبقه‌بندی می‌شود. بعد از حل مدل در دو مرحله، چون برای مشاهداتی که به ناحیه همپوشانی دو گروه تعلق داشتند دو وزن متفاوت در دو مرحله تعیین شد؛ لذا برای به دست آوردن وزن واحد برای تمامی مشاهدات از تابع:

$$z_j = \sum_{f=1}^h \lambda_f^* z_{fj} \quad (11)$$

استفاده می‌کنیم که از معادله زیر محاسبه می‌شود و در آن λ_f^* وزن تخمینی f امین شاخص مالی به دست آمده از $DEA-DA$ می‌باشد و n تعداد کل داده‌ها در مرحله اول و $\#(\text{overlap})$ تعداد کل داده‌ها در مرحله دوم می‌باشد.

$$z_j = \frac{n}{n + \#(\text{overlap})} \sum_{f=1}^h \lambda_f^* z_{fj} + \frac{\#(\text{overlap})}{n + \#(\text{overlap})} \sum_{f=1}^h \lambda_f^* z_{fj} = \sum_{f=1}^h \left(\frac{n}{n + \#(\text{overlap})} \lambda_f^* + \frac{\#(\text{overlap})}{n + \#(\text{overlap})} \lambda_f^* \right) z_{fj} \quad (12)$$

با مرور ادبیات مربوط به $DEA-DA$ ، دیده می‌شود که علاوه بر پژوهش‌های نظری که به تجزیه و تحلیل و ارایه مدل‌ها با کاربردهای مختلف از این تکنیک پرداخته‌اند، پژوهش‌های کاربردی نیز وجود دارد که با استفاده از این مدل‌ها به طبقه‌بندی و تحلیل تشخیص در دنیای تجربی و عملی پرداخته‌اند و از جمله:

سوئیشی و گوتو [۲۷] یک تابع تشخیص ناپارامتریک بر مبنای تحلیل پوششی داده‌ها را برای ارزیابی ورشکستگی شرکت‌ها در ژاپن مطرح کرده و به کار می‌گیرند و نتایج این تابع تشخیص ناپارامتری را با آزمون-های آماری لجیت و پروبیت مقایسه می‌کنند.

همچنین سوئیشی و گوتو [۱۴] در سال ۲۰۱۳ به تحلیل تشخیص ناپارامتریک برای اندازه‌گیری اهمیت هزینه-های تحقیق و توسعه در بورس اوراق بهادار ژاپن پرداختند و نتایج تحقیق نشان‌دهنده اهمیت بالای هزینه‌های تحقیق و توسعه در شرکت‌های فناوری اطلاعات نسبت به دیگر شرکت‌ها بود.

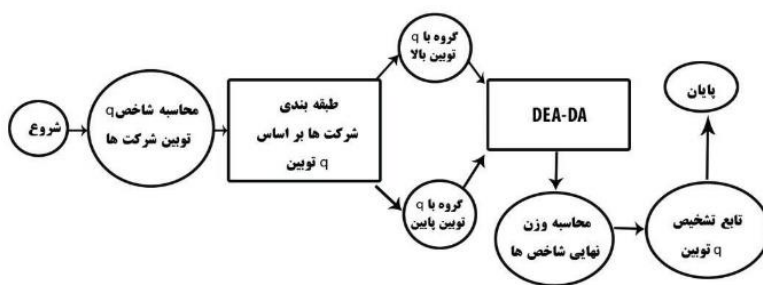
اما از آنجا که تحلیل تشخیص پوششی، تکنیکی جدید می‌باشد؛ لذا پژوهش‌های داخلی اندکی یافت شد که از این تکنیک استفاده کرده باشند که در ادامه به آن‌ها اشاره خواهد شد.

صحت [۲۸] با استفاده از $DEA-DA$ به تجزیه و تحلیل مقایسه‌ای ورشکستگی شرکت‌های مواد غذایی بر مبنای دو مدل جمعی تحلیل پوششی داده‌ها و تحلیل تشخیص پوششی در بین شرکت‌های پذیرفته شده در بورس اوراق بهادار تهران، می‌پردازد و نشان می‌دهد که دقت مدل‌های $DEA-DA$ بیش‌تر از مدل‌های جمعی تحلیل پوششی داده‌ها بوده است.

بیاتی [۲۹] در پژوهشی تحت عنوان "به کارگیری تحلیل پوششی داده‌ها در تحلیل ممیز و کاربرد آن در طبقه‌بندی مشتریان اعتباری بانک ملت" به ارایه مدلی می‌پردازد که با استفاده از قضاوت خبرگان، ویژگی‌های اثرگذار مشتری در سنجش ریسک اعتباری را انتخاب کند و با استفاده از به کارگیری تحلیل پوششی داده‌ها در تحلیل تشخیص، مشتریان خوش حساب را از بد حساب تفکیک نماید.

۳ روش شناسی پژوهش

در این پژوهش، بعد از بررسی داده‌های قابل دسترس، از آنجا که ارزش جایگزینی دارایی‌های شرکت‌ها در بورس اوراق بهادار تهران قابل محاسبه نبود، مدل q توین ساده، رابطه (۱)، به عنوان مدل مناسب پژوهش انتخاب شده و بعد از محاسبه این شاخص و اندازه‌گیری مقدار این شاخص برای تمامی شرکت‌های مورد مطالعه پژوهش، شرکت‌ها را به ترتیب صعودی براساس مقدار q توین آن‌ها مرتب می‌کنیم. سپس برای طبقه‌بندی کردن این شرکت‌ها و یافتن تابع تشخیص مناسب، برای تعیین عضویت هر شرکت جدید، هریک از تکنیک‌های رگرسیون لجستیک، $DEA-DA$ را به ترتیب به کمک نرم افزارهای SPSS و GAMS اجرا می‌کنیم. به عنوان مثال شمای کلی از جریان عملیاتی تکنیک $DEA-DA$ برای به دست آوردن تابع تشخیص q توین همانند شکل ۲ می‌باشد.



شکل ۲. جریان عملیاتی محاسبه تابع تشخیص q توین با $DEA-DA$

این پژوهش به دنبال آن است که از اطلاعات مالی شرکت‌هایی که q خیلی بالا و خیلی پایین دارند استفاده کرده و ارتباط شاخص q را با تعدادی از نسبت‌های مالی پرکاربرد دیگر که پیش از این به عنوان متغیرهای مستقل پژوهش معرفی شدند بیابد و بر همین اساس تابع تشخیص q توین را برای شرکت‌های پذیرفته

شده در بورس اوراق بهادار تهران به کمک روش‌های رگرسیون لجستیک و تحلیل تشخیص پوششی ارایه کند (رابطه‌های (۱۳) - (۱۶)) تا برای تعیین عضویت هر شرکت جدید از این تابع تشخیص استفاده کرده و تعلق آن را به یکی از این دو گروه از قبل تعیین شده، مشخص نماید.

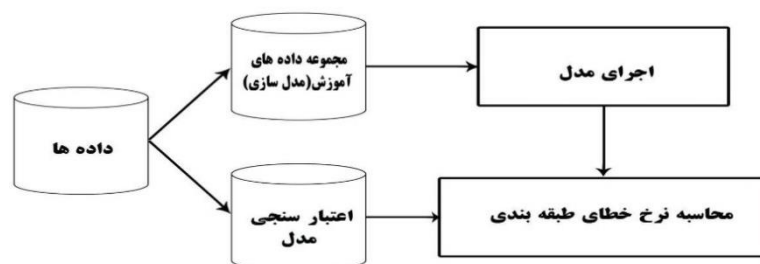
سپس در مرحله بعد، برای بررسی عملکرد مدل‌های به کار گرفته شده مجموعه شرکت‌های مورد آزمون را به دو دسته آموزشی و آزمایشی تقسیم کرده که توسط مجموعه داده آموزشی مدل دسته‌بندی ساخته شده و با مجموعه داده آزمایشی مورد ارزیابی قرار می‌گیرد (شکل ۳).

این پژوهش، داده‌ها را بدین صورت به دو گروه آموزشی و آزمایشی تقسیم می‌کند که دقیقاً ۵۰ درصد داده‌ها را برای مرحله مدل‌سازی (آموزشی) که همان به دست آوردن تابع تشخیص است و ۵۰ درصد باقیمانده را برای اعتبارسنجی مدل به کار می‌گیرد. برای این کار می‌توان به یکی از دو صورت زیر عمل کرد:

۱- ۵۰ درصد اولیه که برای آموزش به کار می‌رود شامل ۲۵ درصد اولیه و ۲۵ درصد نهایی از مجموعه شرکت‌های مورد بررسی است که به ترتیب صعودی براساس مقدار شاخص توین مرتب شده‌است و ۵۰ درصد میانی این داده‌ها به عنوان داده‌های آزمایشی برای بررسی اعتبارسنجی مدل به کار رود.

۲- ۵۰ درصد داده‌ها که برای آموزش در نظر گرفته می‌شود شامل ۵۰ درصد میانی این داده‌ها و مجموعه ۲۵ درصد اولیه و ۲۵ درصد نهایی داده‌ها به عنوان داده‌های آزمایشی برای بررسی اعتبارسنجی مدل به کار رود.

که در این پژوهش از حالت اول استفاده شده است؛ داده‌های ۲۵ درصد بالا و ۲۵ درصد پایین از مجموعه ۱۸۴ مورد بررسی نشان‌دهنده مجموعه شرکت‌هایی هستند که q خیلی بالا و یا q خیلی پایین دارند و می‌توان ادعا کرد که علی‌رغم مورد تردید بودن دقت برآورد شاخص توین، مقدار این شاخص برای این مجموعه داده‌ها قابل اطمینان‌تر باشد و لذا ۵۰ درصد میانی که شاید تغییر خیلی کمی در آن باعث تغییر عضویت گروهی آن واحد شود برای آزمایش مدل به کار خواهیم برد.



شکل ۳. اعتبارسنجی مدل

از آنجا که در این حالت، مجموعه داده انتخابی برای مرحله مدل‌سازی مقادیری کاملاً تفکیک شده‌است و لذا مدل‌سازی براساس آن‌ها راحت‌تر بوده؛ اما مرحله آزمایش و محاسبه دقت تابع به دست آمده، در واقع یک آزمایش یا امتحان سخت‌گیرانه خواهد بود؛ زیرا وقتی از داده‌هایی برای مدل‌سازی و تعیین ضرایب هر یک از شاخص‌ها استفاده شود که کاملاً از یکدیگر جدا هستند؛ لذا در مرحله مدل‌سازی به راحتی الگوی ارتباطی بین متغیرهای مستقل و متغیر وابسته جهت طبقه‌بندی کردن واحدها شناسایی می‌شود و بنابراین استفاده از این ضرایب

نتایج بهتری را به دنبال خواهد داشت؛ لذا بعد به دست آوردن تابع تشخیص در مرحله اول، مقدار تابع تشخیص داده های ۵۰ درصد میانی لیست مرتب شده براساس مقدار شاخص توپین را محاسبه کرده و عضویت تمامی این واحدها را به دست آورده و آنگاه با عضویت گروهی اولیه آن ها مقایسه و میزان خطای طبقه بندی هر یک از تکنیک ها را محاسبه و گزارش می کنیم.

۴ تجزیه و تحلیل نتایج

در ادامه خروجی هر یک از تکنیک های به کار گرفته شده گزارش و تجزیه و تحلیل می شود.

جدول ۱. خروجی رگرسیون لجستیک با ۹ متغیر مستقل

متغیرهای مستقل	ضریب بتا (β)	خطای استاندارد	آماره والد	معنی داری	$\text{Exp}(\beta)$
NITR	۰/۰۴۹	۰/۰۲۰	۵/۸۹۵	۰/۰۱۵	۱/۰۵۰
DPR	۳/۱۷۰	۱/۳۵۱	۵/۵۰۳	۰/۰۱۹	۲۳/۸۰۸
TDTA	-۴/۷۶۶	۴/۵۹۶	۱/۰۷۶	۰/۳۰۰	۰/۰۰۹
CLTA	۳/۴۱۸	۵/۴۵۳	۰/۳۹۳	۰/۵۳۱	۳۰/۵۲۳
WCTA	۲/۰۴۷	۲/۸۰۹	۰/۵۳۱	۰/۴۶۶	۷/۷۴۲
Gror ₁	۲/۸۷۱	۱/۶۶۴	۲/۹۷۷	۰/۰۸۴	۱۷/۶۵۶
LISE	-۰/۰۳۳	۰/۰۶۳	۰/۲۷۶	۰/۵۹۹	۰/۹۶۷
CATA	-۳/۹۸۹	۳/۸۵۸	۱/۰۶۹	۰/۳۰۱	۰/۰۱۹
SIZE	-۱/۵۱۸	۰/۵۹۵	۶/۵۱۱	۰/۰۱۱	۰/۲۱۹
مقدار ثابت	۹/۲۴۰	۴/۰۰۳	۵/۳۲۷	۰/۰۲۱	۱۰۳۰۰/۰۲۶

✓ مقدار خطای طبقه بندی در این حالت برابر با ۳۴/۷۸ درصد محاسبه شد.

و براساس داده های جدول بیان شده، تابع تشخیص شاخص توپین براساس ۹ متغیر مستقل با استفاده از رگرسیون لجستیک به صورت زیر می باشد:

$$y = 0.049 \text{ NITR} + 3.170 \text{ DPR} - 4.766 \text{ TDTA} + 3.418 \text{ CATA} + 2.047 \text{ WCTA} + 2.871 \text{ Gror}_1 - 0.033 \text{ LISE} - 3.989 \text{ CATA} - 1.518 \text{ Size} + 9.240 \quad (13)$$

یکی از آماره های مهم این جدول، آماره والد می باشد که مهم ترین آماره برای آزمون معنی داری حضور هر متغیر مستقل در مدل می باشد، که از طریق مقدار Sig نشان می دهد که حضور کدام یک از متغیرهای مستقل در مدل مفید و اثر آن معنی دار است. مشاهده می شود که شاخص های NITR و DPR و Gror₁ و Size در سطح معنی داری کوچک تر از ۰/۱۰ معنی دار هستند. از سوی دیگر ضریب بتای مثبت نشان دهنده ارتباط موثر

مستقیم متغیر مستقل بر روی طبقه‌بندی واحدها در گروه هدف را که همان گروه ۱ است نشان می‌دهد و ضرایب منفی نشان‌دهنده ارتباط معکوس بین متغیر مستقل و وابسته با هدف تعلق در گروه هدف می‌باشند. با اندکی توجه به ستون ضریب بتا، دیده می‌شود وضعیت شاخص‌های TDTA و CLTA و WCTA و CATA بسیار جالب است؛ زیرا اگرچه ضریب معنی‌داری بالایی دارند و در آزمون والد رد می‌شوند؛ اما مقدار بالای ضریب بتای این شاخص‌ها نشان‌دهنده تاثیر بالای این شاخص‌ها بر متغیر وابسته را نشان می‌دهد. و علت، پراکندگی زیاد انحراف معیار توزیع نمونه‌گیری یا همان خطای معیار می‌باشد که این امر با ملاحظه ستون S.E. کاملاً واضح است.

جدول ۲. خروجی DEA-DA با ۹ متغیر مستقل

شاخص	وزن مرحله اول	وزن مرحله دوم	وزن نهایی
NITR	۰/۰۰۰۶	۰/۰۰۰۳	۰/۰۰۰۴
DPR	۰/۰۶۹	۰/۰۱۸	۰/۰۳۸
TDTA	-۰/۳۶۸	-۰/۲۵۸	-۰/۳۰۱
CLTA	۰/۰۰۰۱	۰/۴۲۵	۰/۲۶۱
WCTA	۰/۲۷۷	۰/۱۴۵	-۰/۰۱۸
Gror1	-۰/۰۱۳	۰/۰۲۲	۰/۰۰۹
LISE	۰/۰۰۸	-۰/۰۰۰۷	۰/۰۰۳
CATA	۰/۲۶۳	-۰/۱۱۴	۰/۰۳۲
SIZE	۰/۰۰۰۱	۰/۰۱۶	-۰/۰۱۰

✓ قابل ذکر است که این نتایج به ازای $M = 10000$ و $\varepsilon = 0/0001$ محاسبه شده‌است و خطای طبقه‌بندی در این حالت ۳۹/۱ درصد بوده است.

و براساس داده‌های این جدول، تابع تشخیص شاخص توبین براساس ۹ متغیر مستقل با استفاده از DEA-DA به صورت زیر می‌باشد:

$$y = 0/0004 \text{ NITR} + 0/038 \text{ DPR} - 0/301 \text{ TDTA} + 0/261 \text{ CATA} - 0/018 \text{ WCTA} + 0/009 \text{ Gror1} + 0/003 \text{ LISE} + 0/032 \text{ CATA} - 0/010 \text{ Size} \quad (14)$$

با اندکی دقت در مقدار وزن شاخص‌ها در هر مرحله مشاهده می‌شود شاخص‌هایی مانند CLTA و WCTA و CATA وضعیت کاملاً بی‌ثباتی را در دو مرحله نشان می‌دهند. مقداری از این بی‌ثباتی را می‌توان بدین ترتیب توجیه کرد که این شاخص‌ها در تامین هدف هر مرحله، نقش متفاوتی را ایفا می‌کنند؛ اما مقدار زیاد این تفاوت در بعضی از این شاخص‌ها باعث شده است که مرحله نهایی، نشان‌دهنده غیر موثر بودن آن‌ها در طبقه‌بندی واحدها باشد.

با نگاهی به ستون Sig و ستون ضریب B در جدول ۱ مشاهده می‌کنیم تعدادی از شاخص‌ها هستند که ضریب بتای آن‌ها بالا بوده؛ اما به علت بالا بودن خطای معیار، این شاخص‌ها دارای Sig غیر معنی‌دار می‌باشند؛ لذا بعد از مقایسه نتایج لجستیک و DEA-DA و مشاهده اینکه تعدادی از این شاخص‌ها (CLTA و WCTA و CATA) در تکنیک DEA-DA نیز بی‌ثباتی در مقدار وزن در دو مرحله آن دیده می‌شود، در ادامه جهت بالا بردن اعتبار نتایج پژوهش بعد از حذف این سه شاخص و همچنین NITR که وزن بسیار کوچکی در هر دو مرحله DEA-DA اختیار کرده است هر دو مدل را دوباره اجرا کرده و نتایج به صورت زیر گزارش می‌گردد:

جدول ۳. خروجی رگرسیون لجستیک با ۵ متغیر مستقل

متغیرهای مستقل	ضریب بتا (β)	خطای استاندارد	آماره والد	معنی‌داری	$\text{Exp}(\beta)$
DPR	۴/۶۲۶	۱/۲۴۳	۱۳/۸۵۵	۰/۰۰	۱۰۲/۰۷۷
TDTA	-۷/۷۳۲	۲/۰۷۲	۱۲/۵۰۸	۰/۰۰	۰/۰۰۱
Gror1	۳/۴۸۸	۱/۴۴۴	۵/۸۳۶	۰/۰۱۶	۳۲/۷۱۸
LISE	-۰/۰۱۹	۰/۰۶۰	۰/۱۰۲	۰/۷۵۰	۰/۹۱۸
SIZE	-۱/۰۳۵	۰/۵۰۱	۴/۲۶۴	۰/۰۳۹	۰/۳۵۵
مقدار ثابت	۷/۴۳۷	۳/۳۷۹	۴/۸۴۵	۰/۰۲۸	۱۶۹۸/۸۵

✓ مقدار خطای طبقه‌بندی در این حالت برابر با ۳۳/۶۰ درصد محاسبه شد.

و براساس داده‌های این جدول، تابع تشخیص شاخص توپین براساس ۵ متغیر مستقل با استفاده از رگرسیون لجستیک به صورت زیر می‌باشد:

$$y = 4/626 \text{ DPR} - 7/732 \text{ TDTA} + 3/488 \text{ Gror1} - 0/019 \text{ LISE} - 1/035 \text{ Size} + 7/437 \quad (15)$$

با توجه به مقدار Sig مربوط به شاخص‌های DPR و TDTA و Gror1 و Size مشاهده می‌شود که این شاخص‌ها در سطح کم‌تر از ۰/۱ دارای ضریب معنی‌داری می‌باشند و این نتیجه با نتایج مربوط ستون بتا نیز همخوانی دارد؛ زیرا همین شاخص‌ها در ستون ضریب بتا نیز مقادیر بالایی را اختیار کرده که نشان از تاثیر زیاد آن‌ها می‌باشد.

مشابه رگرسیون لجستیک، بعد از حذف متغیرهایی که پیش از این شرح داده شد، با کمک نرم افزار GAMS مدل دو مرحله‌ای برنامه ریزی خطی DEA-DA را برای متغیرهای باقیمانده اجرا کرده و نتایج زیر گزارش می‌گردد.

با نگاهی کلی به جدول ۴ به راحتی دیده می شود که باز هم با توجه به تغییر تابع هدف مدل در دو مرحله DEA-DA مقدار هر یک از متغیرهای مستقل اندکی تغییر می کند و در واقع اهمیت آن در تامین هدف هر مرحله کمی متفاوت است؛ اما ثبات کلی این شاخص ها به وضوح دیده می شود؛ زیرا برخلاف حالت قبلی که با ۹ شاخص حل شد، دیگر شاخصی یافت نمی شود که مثلاً از یک مقدار منفی به یک مقدار مثبت بزرگ و بالعکس تغییر پیدا کند.

جدول ۴. خروجی DEA-DA با ۵ متغیر مستقل

شاخص	وزن مرحله اول	وزن مرحله دوم	وزن نهایی
DPR	۰/۳۲۶	۰/۱۸۰	۰/۲۳۴
TDTA	-۰/۳۵۱	-۰/۴۹۴	-۰/۴۴۱
Gror1	۰/۲۰۳	۰/۲۱۲	۰/۲۰۸
LISE	۰/۰۱۰	۰/۰۲۴	۰/۰۱۹
SIZE	-۰/۱۱	-۰/۰۹	-۰/۰۹۷

✓ قابل ذکر است که این نتایج به ازای $M = 100$ و $\varepsilon = 0.001$ محاسبه شده اند و خطای طبقه بندی در این حالت ۳۴/۸ درصد بوده است.

و براساس داده های جدول، تابع تشخیص شاخص توپین براساس ۵ متغیر مستقل با استفاده از DEA-DA به صورت زیر می باشد:

$$y = 0.234 \text{ DPR} - 0.441 \text{ TDTA} + 0.208 \text{ Gror}_1 + 0.019 \text{ LISE} - 0.097 \text{ Size} \quad (16)$$

۵ نتیجه گیری

با نگاهی به خطای طبقه بندی محاسبه شده هر یک از تکنیک های به کار گرفته شده در این پژوهش، در دو حالت مختلف ۹ متغیره و ۵ متغیره، دیده می شود که خطای محاسبه دو روش به هم نزدیک می باشد؛ اما خطای طبقه بندی رگرسیون لجستیک کم تر از تحلیل تشخیص پوششی می باشد. همچنین دیده می شود که شاخص های نسبت کل بدهی ها به کل دارایی ها (TDTA)، نسبت پرداختی سود نقدی (DPR) و نرخ رشد فروش یک ساله (Gror_1) مهم ترین شاخص ها در تعیین تابع تشخیص q توپین می باشند.

خطای طبقه بندی محاسبه شده در تکنیک $DEA-DA$ اندکی بیش تر از رگرسیون لجستیک است؛ اما با توجه به مبانی قدرتمند و مستدل ریاضی و از سوی دیگر با توجه به نو بودن این تکنیک و اینکه هر علمی بعد از دوره معرفی، دوره رشد را در پی داشته و به بلوغ می رسد، انتظار می رود که این تکنیک نیز در سال های آتی بر معایب مدل ها که در پژوهش های کاربردی دیده می شود غلبه کرده و بتواند بسیار قوی تر از این ظاهر شود.

این تحقیق با اجرای مدل دو مرحله ای DEA-DA که در رابطه‌های (۶)–(۱۲) مشخص شده است، از یک طرف به تشخیص شاخص‌های مهم که در تشخیص q توپین اهمیت دارند کمک کرده و از طرف دیگر به بررسی عملکرد و گزارش روش ریاضی DEA-DA می‌پردازد که برای شناسایی و رفع معایب آن نیازمند ارایه گزارش در زمینه‌های عملی است.

ضمناً از این تابع تشخیص به عنوان یک راه برای تشخیص q توپین شرکت‌ها و قضاوت در مورد عملکرد شرکت در کنار دیگر شاخص‌های عملکردی استفاده می‌شود، بخصوص در مورد شرکت‌هایی که قیمت بازار آن‌ها در دسترس نباشد؛ زیرا تمامی متغیرهای پیش‌بین استفاده شده در این پژوهش مستقل از قیمت روز بازار شرکت هستند.

منابع

- [۱] عبدالله زاده درودی، حسن (۱۳۸۸)، «مقایسه محتوای اطلاعاتی نسبت جریان‌های نقدی، ارزش افزوده نقدی، سود حسابداری و ارزش افزوده بازار به ارزش دفتری کل دارایی‌های شرکت در ارزیابی عملکرد شرکت»، پایان نامه کارشناسی ارشد، دانشگاه آزاد اسلامی واحد نیشابور.
- [۲] کیانی نژاد، آزاده (۱۳۸۹)، «بررسی معیارهای عملکرد مالی در شرکت‌های پذیرفته شده در بورس اوراق بهادار تهران»، پایان نامه کارشناسی ارشد مدیریت مالی، دانشگاه آزاد اسلامی نیشابور.
- [۳] شریعت پناهی، مجید (۱۳۸۰)، «اثر نوع مالکیت بر عملکرد مدیران شرکت‌های پذیرفته شده در بورس اوراق بهادار تهران»، رساله دکتری رشته حسابداری، دانشگاه علامه طباطبائی.
- [۵] زراعتگری، رامین (۱۳۸۶)، «بررسی کاربرد q توپین و مقایسه آن با سایر معیارهای ارزیابی عملکرد مدیران در شرکت‌های پذیرفته شده در بورس اوراق بهادار تهران»، پایان نامه کارشناسی ارشد حسابداری، دانشگاه شیراز.
- [۱۶] ستایش، محمدحسین، کارگرفرد جهرمی، محدثه (۱۳۹۰)، «بررسی تاثیر رقابت در بازار محصول بر ساختار سرمایه»، فصلنامه پژوهش‌های تجربی حسابداری مالی، سال اول، شماره اول.
- [۱۷] پورحیدری، امید، هوشمند زعفرانی، رحمت اله، سروستانی، امیر (۱۳۹۱)، «بررسی تأثیر راهبردهای سرمایه در گردش بر سودآوری و ارزش شرکت»، چشم انداز مدیریت مالی، شماره ۷، ۵۵–۷۷.
- [۱۸] نوو، ریموند پی (۱۳۷۶)، مدیریت مالی (ترجمه و اقتباس جهانخانی، علی؛ پارسائیان، علی)، تهران: انتشارات سمت.
- [۱۹] مکیان، سیدنظام الدین، المدرسی، سید محمد تقی، کریمی، تکلوسلیم (۱۳۸۷)، «مقایسه بین مدل شبکه عصبی مصنوعی با روش‌های رگرسیون لجستیک و تحلیل تشخیص آماری در پیش‌بینی ورشکستگی»، فصلنامه پژوهش‌های اقتصادی، سال دهم، شماره دوم، ۱۴۱–۱۶.
- [۲۲] جهانشاهلو، غلامرضا، حسین‌زاده لطفی، فرهاد (۱۳۸۵)، مقدمه‌ای بر تحلیل پوششی داده‌ها، جلد اول، جزوه درسی چاپ نشده، دانشکده ریاضی دانشگاه تربیت معلم.

[۲۸] صحت، سعید (۱۳۹۳)، «تجزیه و تحلیل مقایسه‌ای ورشکستگی شرکت‌های مواد غذایی بر مبنای دو مدل افزایشی تحلیل پوششی داده‌ها (DEA-Additive) و تشخیصی تحلیل پوششی داده‌ها (DEA-DA)»، فصلنامه مطالعات تجربی حسابداری مالی، سال یازدهم، شماره ۴۳، ۱۵۳-۱۸۴.

[۲۹] بیاتی، غلامرضا (۱۳۹۳)، «به کارگیری تحلیل پوششی داده‌ها در تحلیل ممیز و کاربرد آن در طبقه‌بندی مشتریان اعتباری بانک ملت»، پایان نامه کارشناسی ارشد مدیریت صنعتی، دانشگاه آزاد اسلامی واحد تهران مرکزی.

- [4] Wolfe, J., (2003). The tobin's q as a company performance indicator. developments in business simulation and experimental learning 30,155-160.
- [6] Mcconnell, J. J., Servaes, H., (1990). Additional evidence on equity ownership and corporate value. Journal of financial economics, 27(2), 595-612.
- [7] Nabavand, B. Javad, R. (2015). Review between tobin's q with performance evaluation scale based accounting and marketing information in accepted companies in tehran stock exchange. Journal of applied environmental and biological sciences. 4(5s)138-146.
- [8] Heidarpour, F., Malekpoor, S., (2012). Survey the effective factors on tobin's index in tehran stock exchange. World applied sciences journal, 18(4), 575-579..
- [9] Warner, R. M., (2008). Applied statistics: from bivariate through multivariate techniques. Sage.
- [10] Lindenberg, E. B., Ross, S. A., (1981). Tobin's q ratio and industrial organization. Journal of business, 1-32.
- [11] Chung, K. H., Pruitt, S. W., (1994). A simple approximation of tobin's q. financial management, 70-74.
- [13] Steven, B. P., Kenneth, W. W., (1994). Alternative construction of tobins q: an empirical comparison. Journal of empirical financial, 1, 313-341.
- [14] Sueyoshi, T., Goto, M., (2013). A use of dea-da to measure importance of r&d expenditure in japanese information technology industry. Decision support systems, 54(2), 941-952.
- [15] Jensen, M. C., (1997). Eclipse of the public corporation. Harvard business review (sept.-oct. 1989), revised.
- [20] Kao, C., Hwang, S. N., (2010). Efficiency measurement for network systems: it impact on firm performance. Decision support systems, 48(3), 437-446.
- [21] Charnes, A., Cooper, W. W., Rhodes, E., (1978). Measuring the efficiency of decision making units. European journal of operational research, 2(6), 429-444.
- [23] Cooper, W. W., Seiford, L. M., Tone, K., (2002), Data envelopment analysis: a comprehensive text with models, applications. references and dea-solver software. 2nd edition. new york: Springer.
- [24] Sueyoshi, T. (1999). DEA-discriminant analysis in the view of goal programming. European journal of operational research, 115(3), 564-582.
- [25] Sueyoshi, T. (2001). Extended DEA-discriminant analysis. European journal of operational research, 131(2), 324-351.
- [26] Sueyoshi, T. (2004). Mixed integer programming approach of extended dea-discriminant analysis. European journal of operational research, 152(1), 45-55.
- [27] Sueyoshi, T., Goto, M. (2009). Methodological comparison between DEA (data envelopment analysis) and dea-da (discriminant analysis) from the perspective of bankruptcy assessment. European journal of operational research, 199(2), 561-575.