

بررسی دقت مدل های پیش بینی ورشکستگی (مدل های آلتمن، شیراتا، اهلسون، زمیسکی، اسپرینگیت، سی ای اسکور، فولمر، ژنتیک فرج زاده و ژنتیک مک کی) در بورس اوراق بهادار تهران

**حسن قدرتی - استادیار گروه حسابداری دانشگاه آزاد اسلامی واحد کاشان
امیرهادی معنوی مقدم - کارشناس ارشد حسابداری**

چکیده

یکی از روش های پیش بینی تداوم فعالیت شرکتها، استفاده از الگوهای پیش بینی بحران مالی است. در این راستا پژوهش حاضر به بررسی درجه کارآمدی مدل های آلتمن، شیراتا، اهلسون، زمیسکی، اسپرینگیت، سی ای اسکور، فولمر، ژنتیک فرج زاده و ژنتیک مک کی در واقعی بودن نتایج پیش بینی و مقایسه کارآمدی و نتایج پیش بینی مدلها با یکدیگر و همچنین تعدیل ضرایب و تعیین قدرت پیش بینی ورشکستگی آنها در شرکت های پذیرفته شده در بورس اوراق بهادار تهران می پردازد.

تحقیق حاضر مطالعه ای کاربردی است. در این تحقیق از روش استنتاج تحلیلی (استقرایی) و طرح تحقیق پس رویدادی (توصیفی- تحلیلی مبتنی بر تجارب گذشته) استفاده شد. نتایج حاصل از آزمون فرضیه اول نشان داد که الگوهای پیش بینی بحران مالی زمیسکی، اسپرینگیت، سی ای اسکور، ژنتیک فرج زاده و ژنتیک مک کی توانایی پیش بینی تداوم فعالیت شرکت های پذیرفته شده در بورس اوراق بها دار تهران را دارند. فرضیه دوم نیز مورد تأیید قرار گرفت به این ترتیب که مدلهایی که با استفاده از تکنیک های هوش مصنوعی (الگوریتم ژنتیک) مدل سازی شده بودند، نسبت به مدلهایی که با استفاده از تکنیک های آماری مدل سازی شده بودند (مدل های کلاسیک) ، در پیش بینی ورشکستگی از قابلیت بیشتری بر خودار بودند.

کلید واژه ها

ورشکستگی، مدل های پیش بینی بحران مالی، مدل ژنتیک

مقدمه

ورشکستگی واژه ای آشنا برای عموم می باشد. ورشکستگی ممکن است در یک مغازه خرده فروشی کوچک که قادر به پرداخت اجاره نیست و به همین علت بسته می شود و یا در یک شرکت تولیدی بزرگ به دلیل نداشتن نقدینگی مطلوب و زیان مستمر سالانه رخ دهد. سرمایه گذاران و اعتبار دهندگان تمایل زیادی برای پیش بینی ورشکستگی بنگاه ها دارند زیرا در صورت ورشکستگی هزینه های زیادی به آنها تحمیل می شود. پیش بینی تداوم فعالیت واحدهای اقتصادی در دوره های آتی یکی از عناصر مهم در تصمیم گیری جهت سرمایه گذاری می باشد. در طول سالیان متمادی این موضوع مورد توجه پژوهشگران بسیاری واقع شده و الگوهایی جهت پیش بینی ورشکستگی ارائه شده اند. هدف این پژوهش بررسی کارایی ۹ مورد از این الگوها می باشد.

تشریح و بیان موضوع

با توسعه بازارهای مالی و متعاقب آن حاکم شدن وضعیت رقابتی، بسیاری از شرکتها ورشکست و از گردونه رقابت خارج می شوند. این امر موجبات نگرانی صاحبان سرمایه و سایر ذینفعان را فراهم می آورد و آنان برای اینکه از سوخت شدن اصل و فرع سرمایه خود جلوگیری کنند به دنبال روشهایی هستند که بحران های مالی را پیش بینی کنند.

یکی از راه هایی که با استفاده از آن می توان اقدام به بهره گیری مناسب از فرصت های سرمایه گذاری و همچنین جلوگیری از به هدر رفتن منابع کرد، پیش بینی ورشکستگی است. به این ترتیب که اولاً با ارائه هشدار های لازم می توان شرکتها را نسبت به وقوع ورشکستگی هوشیار کرد تا آنها با توجه به این هشدار ها دست به اقدامات مقتضی بزنند و دوم اینکه سرمایه گذاران فرصت های مطلوب سرمایه گذاری را از فرصت های نامطلوب تشخیص دهند و منابعشان را در

فرصت ها و مکان های مناسب سرمایه گذاری کنند. به هر حال نشانه های پریشانی مالی خود را به سرعت نشان نمی دهند بلکه در میان حجم انبوهی از اطلاعات مالی و غیر مالی خود را مستتر می سازند. رمز موفقیت در این زمینه شناسایی به هنگام مشکلات مالی است. این مدلها همانند زنگ خطری مشکلات نهفته در ساختار مالی را آشکار می کنند و امکان عکس العمل به موقع را برای مدیران، سرمایه گذاران و سایر افراد و مراجع ذینفع فراهم می آورند. الگوهای پیش بینی ورشکستگی توابعی هستند که با استفاده از نسبت های مالی، تداوم یا توقف فعالیت واحد تجاری را پیش بینی می کنند. در این پژوهش تلاش می شود کارآمدی برخی از این مدلها- با توجه به شیوه مدل سازی آنها- در بورس اوراق بهادار تهران، مورد بررسی و آزمون قرار گیرد.

ضرورت انجام تحقیق

این روزها شاهد وقوع بحران های تجاری در سراسر جهان هستیم. شدت گرفتن رقابت در عرصه صنایع باعث شده است بسیاری از شرکتها ورشکست شده و از گردونه رقابت خارج شوند. این امر موجبات نگرانی صاحبان سرمایه و به طور کلی جامعه را فراهم آورده است. هشدار اولیه از احتمال ورشکستگی مدیریت و سرمایه گذاران را قادر می سازد تا دست به اقدامات پیش گیرانه بزنند. بستانکاران به شدت نسبت به ریسک سوخت شدن اصل و فرع اعتبارات اعطا شده به مشتریان بالقوه و مشتریان فعلی خود حساس هستند. از آنجایی که ورشکستگی هزینه های اقتصادی و اجتماعی سنگینی را بر جامعه تحمیل می کند از دیدگاه کلان نیز مورد توجه قرار می گیرد. زیرا منابع تلف شده در یک واحد اقتصادی دچار بحران مالی می توانست به فرصت های سود آور دیگری تخصیص یابد. با توجه به این مطالب کلیه افراد و مراجع ذینفع نسبت به پیش بینی ورشکستگی، قبل از وقوع آن علاقه مند هستند.

اهداف اساسی تحقیق

۱-تعدیل ضرایب و تعیین قدرت پیش بینی ورشکستگی مدل های آلتمن، شیراتا، اهلسون، زمیسکی، اسپرینگیت، سی ای اسکور، فولمر، ژنتیک فرج زاده، ژنتیک مک کی در شرکت های پذیرفته شده در بورس اوراق بهادار تهران. ۲-بررسی درجه کارآمدی هر یک از مدلها در واقعی بودن نتایج پیش بینی. ۳-مقایسه کارآمدی و نتایج پیش بینی مدلها با یکدیگر

پیشینه تحقیقات

در طول سالیان گذشته کوشش شده بر مبنای اطلاعات حسابداری مالی و نسبت ها مالی و با استفاده از تکنیک ها، متد ها و متدولوژی های مختلف، مدل های پیش بینی ورشکستگی ایجاد شوند. ارائه مدل های پیش بینی ورشکستگی با تحقیقات بیور در سال ۱۹۶۶ آغاز شد. در این قسمت مروری بر این تحقیقات و همچنین تکنیک های مورد استفاده در ساخت مدل های پیش بینی ورشکستگی خواهیم داشت.

آلتمن (۱۹۶۸): او در مدل Z اسکور از داده های مالی دوره های قبل از ورشکستگی بعنوان متغیر مستقل استفاده کرده و شرکتهای ورشکسته یا سالم به عنوان متغیر های وابسته در نظر گرفته شده اند. این مدل به دقت پیش بینی ۹۵٪ در سال اول قبل از بحران مالی دست یافت. وی در سال ۱۹۸۳ یک اصلاحیه روی مدل انجام و مدل جدیدی به نام 'Z' ارائه کرد.

اسپرینگیت (۱۹۷۸): وی همانند آلتمن از تجزیه و تحلیل ممیزی برای انتخاب ۴ نسبت مالی از میان ۱۹ نسبت که بهترین نسبت ها برای تشخیص شرکت های سالم و ورشکسته بود، استفاده کرد. اسپرینگیت با استفاده از ۴۰ شرکت تولیدی این مدل را آزمون کرد و با نرخ اطمینان ۹۲/۵٪ مواجه شد.

اهلسون (۱۹۸۰): اهلسون از نسبت های مالی و تجزیه و تحلیل لجیت چند بعدی برای ایجاد مدل خود استفاده کرد. مدل وی متشکل از ۹ متغیر روی یک نمونه شامل ۱۰۵ شرکت مواجه با بحران مالی و ۲۰۵۸ شرکت فاقد بحران مالی امتحان گردید که نرخ دقتی حدود ۸۵٪ برای یکسال قبل از ورشکستگی حاصل شد.

فولمر (۱۹۸۴): فولمر از تجزیه و تحلیل چند متغیره برای ارزیابی کاربرد ۴۰ نسبت مالی برای یک نمونه ۶۰ تایی شامل ۳۰ شرکت ورشکسته و ۳۰ شرکت غیرورشکسته استفاده کرد. مدل توانست ۹۶٪ شرکت‌های ورشکسته و ۱۰۰٪ شرکت‌های سالم را به درستی پیش بینی نماید.

زمیسکی (۱۹۸۴): زمیسکی با استفاده از تجزیه و تحلیل پرویت با انتخاب نمونه ای مشتمل بر ۴۰ شرکت دچار بحران مالی و ۸۰ شرکت فاقد بحران مالی به نرخ دقتی حدود ۷۸٪ برای یک سال قبل از ورشکستگی رسید.

CA-SCORE (۱۹۸۷): این مدل با استفاده از روش تجزیه و تحلیل چند متغیره با منظور نمودن ۳۰ نسبت مالی از یک نمونه ۱۷۳ تایی از کارخانجات که فروش معادل ۱ تا ۲۰ میلیون دلار در سال داشتند به وجود آمد که با دقت ۸۳٪ صحت مورد آزمایش قرار گرفته است.

فرج زاده دهکردی (۲۰۰۷): این مدل در صدد پیش‌بینی ورشکستگی شرکت‌های پذیرفته شده در بورس اوراق بهادار تهران با استفاده از تکنیک الگوریتم ژنتیک است. نمونه شامل ۷۲ شرکت ورشکسته و ۷۲ شرکت غیرورشکسته است. مدل توانست ۹۴٪ شرکت‌های نمونه را به صورت صحیح طبقه‌بندی نماید.

تکنیک‌های پیش‌بینی ورشکستگی

تکنیک‌های مورد استفاده در ساخت مدل‌های پیش‌بینی ورشکستگی، در سه گروه طبقه‌بندی می‌شوند: تکنیک‌های آماری، تکنیک‌های هوش مصنوعی و مدل‌های نظری.

تکنیک‌های آماری^۱: تکنیک‌های آماری از ابتدایی‌ترین و رایج‌ترین تکنیک‌ها جهت مدل‌سازی پیش‌بینی ورشکستگی بشمار می‌روند. در این مدل‌ها از روش‌های مدل‌سازی استاندارد کلاسیک استفاده شده است و بر نشانه‌های ناتوانی تجاری تمرکز دارند. متغیرهای مورد استفاده در ساخت این مدل‌ها عموماً اطلاعات مندرج در صورتهای مالی است. مدل‌های آماری خود به دو گروه مدل‌های آماری تک متغیره و چند متغیره تقسیم می‌شوند. تحلیل

تشخیصی، احتمال خطی^۲، لوجیت^۳، پروبیت^۴ و فرایندهای تعدیل ناقص^۵ تشکیل دهنده تکنیک های آماری چند متغیره هستند.

تحلیل تشخیصی چند گانه (MDA): روشی است چند متغیره که پدیده ها را بر اساس ویژگی هایشان به گروه های مانع الجمع طبقه بندی می کند. هدف این روش فراهم آوردن ترکیبی خطی از متغیرهای مستقل (نسبت های مالی) است که بتواند شرکتهای ورشکسته و غیر ورشکسته را به بهترین نحو تفکیک کند. از تحقیقات قابل توجه انجام شده با تکنیک تحلیل تشخیصی چند گانه می توان به آلتمن و همکاران (۱۹۷۷) و فولمر (۱۹۸۹) اشاره کرد.

مدل های لوجیت: این مدل ها بر مبنای یک تابع احتمال تجمعی و با استفاده از نسبت های مالی یک شرکت، احتمال تعلق شرکت به یکی از گروههای از پیش تعیین شده را اندازه گیری می کنند. پس از ۱۹۸۱ و به دلیل محدودیت های موجود در روش های تحلیل تشخیصی چند گانه مطالعات ناتوانی تجاری اغلب بر استفاده از لوجیت تمرکز یافتند. مدل های پروبیت نیز، مشابه با مدلهای لوجیت می باشد. تفاوت اصلی آنها در تابع احتمال ورشکستگی می باشد. به هر حال مدلهای لوجیت نسبت به مدلهای پروبیت از محبوبیت بیشتری برخوردار است، چراکه تحلیل پروبیت در مقایسه با تحلیل لوجیت به دلیل استفاده از برآوردهای غیرخطی، به محاسبات بیشتری نیاز دارد.

تکنیک های هوش مصنوعی^۶ (AIT): تکنیک های هوش مصنوعی، مشابه با هوش و منطق انسان، سیستمی است که یاد می گیرد و عملکرد حل مساله خود را با توجه به تجربیات گذشته بهبود می بخشد. تکنیک های هوش مصنوعی به دلیل کارایی بالا و فارغ بودن از مفروضات محدود کننده موجود در روش های آماری با استقبال زیادی مواجه شده اند. تکنیک های هوشمند از الگوریتم های بازگشتی (درخت های تصمیم)^۷، استلال مبتنی بر موضوع^۸، شبکه های عصبی مصنوعی^۹، الگوریتم ژنتیک و مجموعه های سخت^{۱۰} تشکیل شده اند.

الگوریتم ژنتیک^{۱۱}: ایده اصلی الگوریتم ژنتیک، انتقال خصوصیات موروثی توسط ژنها است. الگوریتم ژنتیک یک روش جستجوی احتمالی است که از شبیه سازی تکامل زیستی و طبیعی استفاده می کند. الگوریتم های ژنتیک با بکارگیری اصل بقای بهترین ها، برای تولید تخمین های هر چه بهتر از یک جواب (کروموزوم ها) ، روی جمعیتی از جواب های بالقوه عمل می نماید.

مدل های نظری^{۱۲}: بر خلاف مدل های آماری و تکنیک های هوش مصنوعی که بر نشانه های ناتوانی تجاری تمرکز دارند، مدل های نظری به دنبال تعیین دلایل کیفی ناتوانی تجاری اند. مدل های نظری از نظر ماهیت چند متغیره بوده و معمولاً از تکنیک های آماری برای پشتیبانی کمی مباحث نظری استفاده می کنند.

فرضیه های تحقیق

۱ . مدل های پیش گفته دارای قدرت پیش بینی قابل قبول توقف فعالیت در شرکت های پذیرفته شده در بورس اوراق بهادار تهران هستند.

۲ . مدل هایی که با استفاده از تکنیک های هوش مصنوعی (الگوریتم ژنتیک) مدل سازی شده اند نسبت به مدلهایی که با استفاده از تکنیک های آماری مدل سازی شده اند (مدل های کلاسیک)، در پیش بینی ورشکستگی از قابلیت بیشتری بر خودار می باشند.

روش تحقیق

تحقیق حاضر مطالعه ای کاربردی است. در این تحقیق از روش استنتاج تحلیلی (استقرایی) و طرح تحقیق پس رویدادی (توصیفی_تحلیلی مبتنی بر تجارب گذشته) استفاده می شود.

ابزار گرد آوری اطلاعات

اطلاعات مورد نیاز از طریق اسنادکاوی، مشاهده، مطالعه صورت های مالی حسابرسی شده، بررسی مستندات و گزارش های منتشر شده توسط سازمان بورس و بانک های اطلاعاتی گردآوری شدند.

روش های تجزیه و تحلیل اطلاعات

برای آزمون فرضیه ها از روش تحلیل رگرسیونی استفاده شده است. در آزمون فرضیه ها از روش آماری Binary Logistic Analysis استفاده می گردد. همچنین از آماره آزمون والد، برای استخراج ضرایب مؤثر با توجه به شرایط مالی و اقتصادی ایران استفاده می شود. برای آزمون نرمال بودن داده ها، از آزمون کلموگروف-اسمیرنف استفاده می شود. داده های جمع آوری شده با استفاده از برنامه Excel طبقه بندی گردیده و سپس متغیرهای مورد نظر برای آزمون فرضیه های تحقیق براساس الگوهای نام برده محاسبه می شوند.

جامعه آماری و نمونه آماری

جامعه آماری این پژوهش شامل شرکت های پذیرفته شده در بورس اوراق بهادار تهران طی دوره پنج ساله ۱۳۸۲ تا ۱۳۸۶ است. شرکت های پذیرفته شده در بورس اوراق بهادار تهران، با در نظر گرفتن ویژگی های زیر انتخاب می شوند: ۱- دوره مالی به ۲۹ اسفند منتهی باشد. ۲- طی دوره تحقیق تغییر سال مالی نداشته باشد. ۳- جزء شرکت های سرمایه گذاری، واسطه گری مالی، هلدینگ، لیزینگ و بانک نباشد. ۴- صورت های مالی شرکت در دسترس باشد.

جامعه آماری این پژوهش به دو گروه شرکت های موفق و شرکت های ناموفق تقسیم می شود. شرکت هایی که سه سال پیاپی مشمول ماده ۱۴۱ قانون تجارت بودند، به عنوان شرکتهای ناموفق انتخاب شدند. طبق ماده ۱۴۱، چنانچه میزان زیان انباشته شرکتی از نصف سرمایه اش بیشتر گردد، این شرکت باید سرمایه خود را کاهش و یا فعالیت خود را متوقف نماید. معیار اصلی انتخاب شرکت های گروه موفق یا دارای تداوم فعالیت شاخص Q توین ساده است:

ارزش بازار سهام عادی و ممتاز + ارزش دفتری بدهی ها

$$Q = \text{توین ساده}$$

ارزش دفتری کل دارایی ها در پایان سال

هنگامی که Q توین بیشتر از یک گردد، این امر بیانگر آن است که انگیزه سرمایه گذاری در این شرکت ها وجود دارد. شرکتهای موفق از میان شرکتهایی که Q توین آن ها سه سال متوالی بیش از یک گردید، با استفاده از روش نمونه گیری تصادفی و به تعداد شرکتهای ورشکسته انتخاب شدند. بدین ترتیب با توجه به معیارهای انتخاب شرکتهای گروه اول و دوم ۳۰ شرکت به عنوان شرکت موفق و ۳۰ شرکت به عنوان شرکت ناموفق انتخاب شدند.

مدل تحقیق

در تحلیل های رگرسیونی لازم است متغیرهای کیفی به متغیرهای کمی تبدیل شوند، لذا با تعریف متغیر تصنعی Z متغیرهای کیفی به متغیرهای کمی تبدیل می شوند:

$$Z = \begin{cases} 0 & \text{اگر شرکت ورشکسته است (توقف فعالیت)} \\ 1 & \text{اگر شرکت سالم است (تداوم فعالیت)} \end{cases}$$

متغیر وابسته تداوم فعالیت و توقف فعالیت و متغیر مستقل نسبت های مالی مدل های پیش بینی ورشکستگی به شرح زیر می باشد:

مدل آلتمن: $Z = \beta_0 + \beta_1 A + \beta_2 B + \beta_3 C + \beta_4 D + \beta_5 E$ که در آن؛ A : نسبت سرمایه در گردش به کل دارایی ها، B : نسبت سود انباشته به کل دارایی ها، C : نسبت سود قبل از بهره و مالیات به کل دارایی ها، D : نسبت ارزش بازار حقوق صاحبان سهام به ارزش دفتری کل بدهی ها، E : نسبت فروش به کل دارایی ها.

مدل شیراتا: $Z = \beta_0 + \beta_1 A + \beta_2 B + \beta_3 C + \beta_4 D$ که در آن؛ A: نسبت سود انباشته به کل دارایی، B: نسبت بدهی ها و حقوق صاحبان سهام سال جاری به بدهی ها و حقوق صاحبان سهام سال قبل، C: نسبت هزینه بهره به میانگین حاصل جمع وامها و بدهی ها و اوراق قرضه و اسناد دریافتی تنزیل شده، D: نسبت میانگین حاصل جمع حسابهای پرداختی و اسناد پرداختی ضرب در ۱۲ به فروش.

مدل اهلسون: $Z = \beta_0 + \beta_1 A + \beta_2 B + \beta_3 C + \beta_4 D + \beta_5 E + \beta_6 F + \beta_7 G + \beta_8 H + \beta_9 I$ که در آن؛ A: لگاریتم (کل دارایی ها به شاخص تولید ناخالص ملی (GNP)) بر مبنای ده؛ B: نسبت کل بدهی ها به کل دارایی ها؛ C: نسبت سرمایه در گردش به کل دارایی ها؛ D: نسبت بدهی جاری به دارایی جاری؛ E: عدد یک اگر بدهی کل بیشتر از دارایی کل شود وگرنه عدد صفر؛ F: نسبت سود خالص به کل دارایی ها؛ G: نسبت وجوه نقد حاصل از عملیات به کل بدهی ها؛ H: عدد یک اگر سود خالص برای دو سال گذشته منفی باشد وگرنه عدد صفر؛ I: نسبت میزان تغییر در سود خالص به مجموع قدر مطلق سود هر دو سال.

مدل زمیسکی: $Z = \beta_0 + \beta_1 A + \beta_2 B + \beta_3 C$ که در آن؛ A: نسبت سود خالص به کل دارایی ها؛ B: نسبت کل بدهی ها به کل دارایی ها؛ C: نسبت دارایی های جاری به بدهی های جاری.

مدل اسپرینگیت: $Z = \beta_0 + \beta_1 A + \beta_2 B + \beta_3 C + \beta_4 D$ که در آن؛ A: نسبت سرمایه در گردش به کل دارایی ها؛ B: نسبت سود خالص قبل از بهره و مالیات به کل دارایی ها؛ C: نسبت سود خالص قبل از مالیات به بدهی های جاری؛ D: نسبت فروش به کل دارایی ها.

مدل سی ای اسکور: $Z = \beta_0 + \beta_1 A + \beta_2 B + \beta_3 C$ که در آن؛ A: نسبت حقوق صاحبان سهام (۱) به کل دارایی (۱)؛ B: نسبت هزینه های مالی بعلاوه سود قبل از مالیات و اقلام غیر مترقبه (۱) به کل دارایی (۱)؛ C: نسبت فروش (۲) به کل دارایی (۲). عدد (۱) رقمهای باقیمانده از یک دوره قبل، عدد (۲) رقمهای باقیمانده از دو دوره قبل.

مدل فولمر: $Z = \beta_0 + \beta_1 A + \beta_2 B + \beta_3 C + \beta_4 D + \beta_5 E + \beta_6 F + \beta_7 G + \beta_8 H + \beta_9 I$ که در آن؛ A: نسبت سود انباشته به کل دارایی؛ B: نسبت فروش به کل دارایی؛ C: نسبت سود قبل از مالیات به حقوق صاحبان سهام؛ D: نسبت جریانات نقدی به کل بدهی؛ E: نسبت بدهی به کل دارایی ها؛ F: نسبت بدهی های جاری به کل دارایی؛ G: لگاریتم کل دارایی های مشهود؛ H: نسبت سرمایه در گردش به کل بدهی؛ I: نسبت لگاریتم سود قبل از بهره و مالیات به بهره. مدل ژنتیک فرج زاده: $((C * -9.66) - E) * D)^3 + A, (((B - D) - A) * 6.2)$ ، $((a - (1 - D)) + (A^2))$ ، $a: (IF(E - 5.5) < (-2.5 * C) \text{ then } 1, \text{else } 0)$ ؛ که در آن؛ A: نسبت دارایی های آتی به کل دارایی ها؛ B: نسبت کل بدهی ها به کل دارایی ها؛ C: نسبت فروش به دارایی های جاری؛ D: نسبت سود عملیاتی به فروش؛ E: نسبت هزینه بهره به سود ناخالص.

مدل ژنتیک مک کی: $(A, B, C) = X^2 / (X^2 + Y^2)$ که در آن؛ $X = ((A + 0.85) B) - 0.85$ ، $Y = (1 + C)$ ، A: لگاریتم کل دارایی ها تقسیم بر هزار بر مبنای ده؛ B: نسبت سود خالص به کل دارایی ها؛ C: نسبت وجه نقد به بدهی های جاری. لازم به یاد آوری است که نسبت های مالی مورد استفاده در این الگو ها از الگوهای اصلی این پژوهشگران استخراج شده است و فقط ضرایب متغیر های مستقل در بورس اوراق بهادار تهران تعدیل شده است.

بررسی نرمال بودن داده ها

برای آزمون نرمال بودن داده ها، از آزمون کلموگروف-اسمیرنف استفاده می شود. فرضیه های آزمون به صورت زیر مطرح می شود:

H_0 : داده ها نرمال است (داده ها از جامعه نرمال آمده اند)

H_1 : داده ها نرمال نیست (داده ها از جامعه نرمال نیامده اند)

از آنجا که برنامه ریزی ژنتیک یک تکنیک ناپارامتریک محسوب می شود، اصولاً نیازی به وجود فرضیه هایی در رابطه با نحوه توزیع نسبت های مالی مانند تبعیت از فرض توزیع نرمال

این نسبت ها وجود ندارد. در مورد مدل های کلاسیک که متکی به برخی فرضیات محدود کننده آماری نظیر توزیع نرمال نسبت های مالی در گروه شرکتهای ورشکسته و غیر ورشکسته است، بررسی نحوه پراکندگی این نسبت ها حائز اهمیت است. نتایج آزمون نرمال بودن داده ها در جدول ۱ خلاصه شده است.

اگر مقدار سطح معنی داری بزرگتر از مقدار خطا باشد فرض صفر را نتیجه می گیریم و در صورتی که مقدار سطح معنی داری کوچکتر از خطا باشد فرض یک را نتیجه می گیریم. چون مقدار سطح معنی داری این متغیرها بزرگتر از مقدار خطا $0/05$ می باشد پس فرض صفر را نتیجه می گیریم یعنی این متغیرها همگی نرمال می باشند و برای آزمودن آنها می توانیم از روشهای پارامتری استفاده نمائیم.

جدول ۱- بررسی نرمال بودن داده ها

متغیر	سطح معنی داری	مقدار خطا	تایید فرضیه	متغیر	سطح معنی داری	تایید فرضیه	نتیجه گیری
A آلتمن	۰/۰۶۹	۰/۰۵	H ₀	B زمیسی	۰/۱۱۹	H ₀	نرمال
B آلتمن	۰/۱۳۶	۰/۰۵	H ₀	C زمیسی	۰/۹۹۱	H ₀	نرمال
C آلتمن	۰/۱۶۹	۰/۰۵	H ₀	A اسپرینگیت	۰/۰۹۸	H ₀	نرمال
D آلتمن	۰/۱۰۲	۰/۰۵	H ₀	B اسپرینگیت	۰/۱۶۹	H ₀	نرمال
E آلتمن	۰/۹۲۲	۰/۰۵	H ₀	C اسپرینگیت	۰/۴۰۷	H ₀	نرمال
A شیراتا	۰/۱۵۷	۰/۰۵	H ₀	D اسپرینگیت	۰/۹۲۲	H ₀	نرمال
B شیراتا	۰/۹۴۷	۰/۰۵	H ₀	A سی ای اسکور	۰/۱۱۵	H ₀	نرمال
C شیراتا	۰/۷۸۸	۰/۰۵	H ₀	B سی ای اسکور	۰/۱۶۹	H ₀	نرمال
D شیراتا	۰/۰۷۳	۰/۰۵	H ₀	C سی ای اسکور	۰/۰۸۶	H ₀	نرمال
A اهلسون	۰/۹۲۳	۰/۰۵	H ₀	A فولمر	۰/۱۶۵	H ₀	نرمال
B اهلسون	۰/۱۲۳	۰/۰۵	H ₀	B فولمر	۰/۹۲۲	H ₀	نرمال
C اهلسون	۰/۱۴۰	۰/۰۵	H ₀	C فولمر	۰/۱۱۱	H ₀	نرمال
D اهلسون	۰/۱۶۱	۰/۰۵	H ₀	D فولمر	۰/۳۷۱	H ₀	نرمال
E اهلسون	۰/۱۳۲	۰/۰۵	H ₀	E فولمر	۰/۰۸۲	H ₀	نرمال
F اهلسون	۰/۰۶۵	۰/۰۵	H ₀	F فولمر	۰/۰۹۰	H ₀	نرمال
G اهلسون	۰/۳۷۱	۰/۰۵	H ₀	G فولمر	۰/۹۸۵	H ₀	نرمال
H اهلسون	۰/۱۵۲	۰/۰۵	H ₀	H فولمر	۰/۹۹۷	H ₀	نرمال
I اهلسون	۰/۳۶۴	۰/۰۵	H ₀	I فولمر	۰/۱۰۲	H ₀	نرمال
A زمیسی	۰/۰۶۵	۰/۰۵	H ₀	-	-	-	نرمال

تعدیل مدل ها و سنجش توانایی آن ها

مدل های آلتمن، شیراتا، فولمر و اهلسون : جداول ۲ و ۳ نتیجه آزمون این چهار مدل را نمایش می دهد.

جدول ۲- نتیجه آزمون مدل های شیراتا و آلتمن

متغیر توضیحی	مدل آلتمن		متغیر توضیحی	مدل شیراتا	
مقدار ثابت	۱۳/۴۳۸	B ₀	مقدار ثابت	۱۷/۴۲۷	B ₀
	۵۹۸۸/۲۸۲	S.E.		۷۴۱۰/۰۴۱	S.E.
	۰/۰۰۰	Wald		۰/۰۰۰	Wald
	۰/۹۹۸	سطح معنی داری		۰/۹۹۸	سطح معنی داری
آلتمن A	۱۵/۸۱۷	B ₁	شیراتا A	۱۲۹/۳۰۲	B ₁
	۱۰۸۷۳/۴۹۰	S.E.		۱۰۵۹۶/۰۲۴	S.E.
	۰/۰۰۰	Wald		۰/۰۰۰	Wald
	۰/۹۹۹	سطح معنی داری		۰/۹۹۰	سطح معنی داری
آلتمن B	۹۱/۰۰۵	B ₂	شیراتا B	۱۵/۸۰۲	B ₂
	۱۰۹۸۲/۰۶۲	S.E.		۱۲۱۱۱/۰۵۸	S.E.
	۰/۰۰۰	Wald		۰/۰۰۰	Wald
	۰/۹۹۳	سطح معنی داری		۰/۹۹۹	سطح معنی داری
آلتمن C	۷۰/۹۷۵	B ₃	شیراتا C	-۱۱۴/۰۶۲	B ₃
	۲۷۲۹۸/۱۰۴	S.E.		۹۹۷۰۸/۶۴۳	S.E.
	۰/۰۰۰	Wald		۰/۰۰۰	Wald
	۰/۹۹۸	سطح معنی داری		۰/۹۹۹	سطح معنی داری
آلتمن D	۸/۷۴۵	B ₄	شیراتا D	-۱/۰۱۸	B ₄
	۵۳۵۴/۳۱۲	S.E.		۴۷۲/۶۱۶	S.E.
	۰/۰۰۰	Wald		۰/۰۰۰	Wald
	۰/۹۹۹	سطح معنی داری		۰/۹۹۸	سطح معنی داری
آلتمن E	-۸/۹۰۹	B ₅	-	-	-
	۱۷۹۵۴/۵۳۷	S.E.		-	-
	۰/۰۰۰	Wald		-	-
	۱/۰۰۰	سطح معنی داری		-	-

به علت اینکه فرضیه های تحقیق در سطح اطمینان ۹۰ درصد آزمون می شوند، متغیرهایی که سطح معنی داری کمتر از ۱۰ درصد دارند، به عنوان متغیرهای موثر که از قدرت توضیح دهندگی خوبی برخوردارند، انتخاب می شوند.

جدول ۳- نتیجه آزمون مدل های فولمر و اهلسون

متغیر توضیحی	مدل اهلسون		متغیر توضیحی	مدل فولمر	
مقدار ثابت	۱۰۹/۹۹۹	B ₀	مقدار ثابت	۰/۴۱۴	B ₀
	۳۲۴۷۵/۱۶۵	S.E.		۶۳۹۶۰/۸۳۸	S.E.
	۰/۰۰۰	Wald		۰/۰۰۰	Wald
	۰/۹۹۷	سطح معنی داری		۱/۰۰۰	سطح معنی داری
اهلسون A	۱۶/۲۶۱	B ₁	فولمر A	۳۰/۶۳۴	B ₁
	۷۱۱۲/۰۵۹	S.E.		۱۸۷۱۵/۹۴۰	S.E.
	۰/۰۰۰	Wald		۰/۰۰۰	Wald
	۰/۹۹۸	سطح معنی داری		۰/۹۹۹	سطح معنی داری
اهلسون B	-۳۶/۰۲۵	B ₂	فولمر B	۱۱/۴۴۱	B ₂
	۲۹۱۳۶/۱۷۷	S.E.		۱۰۷۵۵/۴۵۳	S.E.
	۰/۰۰۰	Wald		۰/۰۰۰	Wald
	۰/۹۹۹	سطح معنی داری		۰/۹۹۹	سطح معنی داری
اهلسون C	-۱۱/۸۹۷	B ₃	فولمر C	-۱/۱۹۷	B ₃
	۳۵۱۱۸/۶۳۴	S.E.		۱۹۸۱/۸۰۵	S.E.
	۰/۰۰۰	Wald		۰/۰۰۰	Wald
	۱/۰۰۰	سطح معنی داری		۱/۰۰۰	سطح معنی داری
اهلسون D	۵/۱۲۹	B ₄	فولمر D	۵۶/۰۰۵	B ₄
	۱۱۵۳۳/۴۰۳	S.E.		۳۷۱۴۰/۲۷۵	S.E.
	۰/۰۰۰	Wald		۰/۰۰۰	Wald
	۱/۰۰۰	سطح معنی داری		۰/۹۹۹	سطح معنی داری
اهلسون E	-۸۵/۶۶۰	B ₅	فولمر E	-۱۲۳/۳۹۵	B ₅
	۲۷۸۱۷/۵۷۳	S.E.		۴۳۶۵۷/۹۲۱	S.E.
	۰/۰۰۰	Wald		۰/۰۰۰	Wald

	۰/۹۹۸	سطح معنی داری		۰/۹۹۸	سطح معنی داری
F اهلسون	-۱۷۳/۴۳۹	B ₆	F فولمر	۷۸/۲۲۱	B ₆
	۵۵۲۸۶/۱۵۷	S.E.		۴۴۳۲۳/۲۴۳	S.E.
	۰/۰۰۰	Wald		۰/۰۰۰	Wald
	۰/۹۹۷	سطح معنی داری		۰/۹۹۹	سطح معنی داری
G اهلسون	۲۷/۸۴۱	B ₇	G فولمر	۸/۶۸۸	B ₇
	۲۰۱۳۸/۵۸۴	S.E.		۱۲۴۹۸/۸۵۱	S.E.
	۰/۰۰۰	Wald		۰/۰۰۰	Wald
	۰/۹۹۹	سطح معنی داری		۰/۹۹۹	سطح معنی داری
H اهلسون	-۹۲/۱۰۷	B ₈	H فولمر	667/15	B ₈
	۲۱۶۹۷/۲۱۱	S.E.		۱۹۶۸۰/۹۱۵	S.E.
	۰/۰۰۰	Wald		۰/۰۰۰	Wald
	۰/۹۹۷	سطح معنی داری		۰/۹۹۹	سطح معنی داری
I اهلسون	۵۸/۷۱۸	B ₉	I فولمر	۲۴۶۷/۸۲۸	B ₉
	۱۰۶۶۲/۹۸۱	S.E.		۶۳۳۴۱۱۵	S.E.
	۰/۰۰۰	Wald		۰/۰۰۰	Wald
	۰/۹۹۶	سطح معنی داری		۱/۰۰۰	سطح معنی داری

در این چهار مدل هیچ متغیر موثری که از توضیح دهندگی خوبی جهت پیش بینی توقف فعالیت شرکتها برخوردار باشد وجود ندارد، زیرا مقدار سطح معنی داری کلیه متغیرهای مستقل بیشتر از مقدار ۰/۱ می باشد. پس این مدلها دارای قدرت پیش بینی توقف فعالیت در شرکت های پذیرفته شده در بورس اوراق بهادار تهران نیستند.

مدل های زمیسکی، اسپرینگیت، سی ای اسکور: جدول ۴ نتیجه آزمون این سه مدل را نمایش می دهد.

جدول ۴- نتیجه آزمون مدل های سی ای اسکور، زمیسی و اسپرینگیت

متغیر توضیحی	مدل زمیسی		متغیر توضیحی	مدل سی ای اسکور	
مقدار ثابت	۱۰/۳۴۹	B ₀	مقدار ثابت	۳/۲۰۴	B ₀
	۷/۸۳۱	S.E.		۲/۴۱۲	S.E.
	۱/۷۴۶	Wald		۱/۷۶۴	Wald
	۰/۱۸۶	سطح معنی داری		۰/۱۸۴	سطح معنی داری
A زمیسی	۱۰/۲۷۵	B ₁	A سی ای اسکور حقوق صاحبان سهام (۱) به کل دارایی (۱)	۱۲/۹۲۳	B ₁
	۱۲/۶۹۵	S.E.		۵/۸۳۵	S.E.
	۰/۶۵۵	Wald		۴/۹۰۶	Wald
	۰/۴۱۸	سطح معنی داری		۰/۰۲۷	سطح معنی داری
B زمیسی نسبت کل بدهی ها به کل دارایی ها	-۱۲/۳۳۷	B ₂	B سی ای اسکور	۹/۷۷۵	B ₂
	۷/۳۶۷	S.E.		۱۰/۱۹۷	S.E.
	۲/۸۰۴	Wald		۰/۹۱۹	Wald
	۰/۰۹۴	سطح معنی داری		۰/۳۳۸	سطح معنی داری
C زمیسی	۰/۳۰۴	B ₃	C سی ای اسکور	۱/۵۹۶	B ₃
	۲/۸۲۰	S.E.		۳/۴۴۷	S.E.
	۰/۰۱۲	Wald		۰/۲۱۴	Wald
	۰/۹۱۴	سطح معنی داری		۰/۶۴۳	سطح معنی داری
متغیر توضیحی	مدل اسپرینگیت		متغیر توضیحی	مدل اسپرینگیت	
مقدار ثابت	-۰/۳۹۰	B ₀	C اسپرینگیت	۱۶/۴۷۱	B ₃
	۱/۸۵۵	S.E.		۱۱/۴۳۸	S.E.
	۰/۰۴۴	Wald		۲/۰۷۴	Wald
	۰/۸۳۴	سطح معنی داری		۰/۱۵۰	سطح معنی داری

A اسپرینگیت نسبت سرمایه در گردش به کل دارایی ها	۸/۵۰۳	B ₁	D اسپرینگیت	۰/۷۴۰	B ₄
	۴/۴۳۷	S.E.		۲/۹۵۶	S.E.
	۳/۶۷۲	Wald		۰/۰۶۳	Wald
	۰/۰۵۵	سطح معنی داری		۰/۸۰۲	سطح معنی داری
B اسپرینگیت	-۰/۵۲۲	B ₂	-	-	-
	۱۷/۱۲۳	S.E.		-	-
	۰/۰۰۱	Wald		-	-
	۰/۹۷۶	سطح معنی داری		-	-

متغیر موثری که از توضیح دهندگی بسیار خوبی جهت پیش بینی توقف فعالیت شرکت ها برخوردار بوده است در مدل زمیسکی، B یعنی نسبت کل بدهی ها به کل دارایی ها؛ در مدل اسپرینگیت، A یعنی نسبت سرمایه در گردش به کل دارایی ها و در مدل سی ای اسکور، A یعنی نسبت حقوق صاحبان سهام به کل دارایی می باشد، زیرا سطح معنی داری آنها در سطح اطمینان ۹۰ درصد کمتر از ۱۰ درصد بوده است. الگوهای تعدیل شده زمیسکی، اسپرینگیت، سی ای اسکور براساس آزمون انجام شده بر مبنای متغیرهای موثر به شرح زیر می باشند:

$$Z = -12.337 B$$

$$Z = 8.503 A$$

$$Z = 12.923 A$$

الگوی تعدیل شده زمیسکی

الگوی تعدیل شده اسپرینگیت

الگوی تعدیل شده سی ای اسکور

جدول ۵- توانایی و دقت الگوهای زمیسی، اسپرینگیت، سی ای اسکور

دقت کلی		پیش بینی شده مدل زمیسی				گروه شرکت ها
درصد	تعداد	موفق		ناموفق		
		درصد	تعداد	درصد	تعداد	
۹۰	۳۰	۱۰	۳	۹۰	۲۷	ناموفق
۹۰	۳۰	۹۰	۲۷	۱۰	۳	موفق
۹۰	۶۰	–	۳۰	–	۳۰	جمل کل
دقت کلی		پیش بینی شده مدل اسپرینگیت				گروه شرکت ها
درصد	تعداد	موفق		ناموفق		
		درصد	تعداد	درصد	تعداد	
۹۰	۳۰	۱۰	۳	۹۰	۲۷	ناموفق
۹۰	۳۰	۹۰	۲۷	۱۰	۳	موفق
۹۰	۶۰	–	۳۰	–	۳۰	جمل کل
دقت کلی		پیش بینی شده مدل سی ای اسکور				گروه شرکت ها
درصد	تعداد	موفق		ناموفق		
		درصد	تعداد	درصد	تعداد	
۹۰	۳۰	۱۰	۳	۹۰	۲۷	ناموفق
۹۰	۳۰	۹۰	۲۷	۱۰	۳	موفق
۹۰	۶۰	–	۳۰	–	۳۰	جمل کل

جدول ۵ توانایی و دقت الگوهای زمیسی، اسپرینگیت، سی ای اسکور را با توجه به اطلاعات شرکتهای نمونه ارائه کرده است. در ستون ناموفق عدد ۳ و درصد ۱۰ نشان دهنده خطای نوع دوم و در ستون موفق نیز عدد ۳ و درصد ۱۰ نشان دهنده خطای نوع اول است. در جدول بالا

هر سه الگوی تعدیل شده دارای دقت کلی ۹۰ درصد می باشد. پس دقت پیش بینی سه مدل تایید شده زمیسی، اسپرینگیت و سی ای اسکور با یکدیگر برابر می باشد.

مدل های ژنتیک فرج زاده و ژنتیک مک کی: مدل های ژنتیک از نظر تکنیک

آماري و شیوه به وجود آمدن با مدل های کلاسیک متفاوت هستند. به همین جهت نیازی به تعدیل ضرایب در آنها وجود ندارد. در ضمن مدل ژنتیک فرج زاده در شرایط محیطی ایران به وجود آمده و نیازی به تعدیل بر اساس شرایط محیطی بورس اوراق بها دار تهران در آن وجود ندارد. از این رو ابتدا نسبت های مورد نظر در مدل محاسبه شده و سپس در سه زیر شاخه مدل قرار داده شده و اعداد بدست آمده با هم جمع می شوند .

جدول ۶- توانایی و دقت الگوهای ژنتیک فرج زاده و ژنتیک مک کی

دقت کلی		پیش بینی شده مدل ژنتیک فرج زاده				گروه شرکت ها
درصد	تعداد	موفق		ناموفق		
		درصد	تعداد	درصد	تعداد	
۹۶/۷	۳۰	۳/۳	۱	۹۶/۷	۲۹	ناموفق
۸۶/۷	۳۰	۸۶/۷	۲۶	۱۳/۳	۴	موفق
۹۱/۷	۶۰	-	۲۷	-	۳۳	جمل کل
دقت کلی		پیش بینی شده مدل ژنتیک مک کی				گروه شرکت ها
درصد	تعداد	موفق		ناموفق		
		درصد	تعداد	درصد	تعداد	
۸۳/۳	۳۰	۱۶/۷	۵	۸۳/۳	۲۵	ناموفق
۱۰۰	۳۰	۱۰۰	۳۰	۰	۰	موفق
۹۱/۶۵	۶۰	-	۳۵	-	۲۵	جمل کل

در مدل مک کی نیز ابتدا نسبت های مورد نظر در مدل محاسبه شده و سپس به وسیله آنها X و Y محاسبه شده و در فرمول مورد نظر قرار می گیرد. در صورتی که عدد بدست آمده بزرگتر از پنج دهم باشد، شرکت به عنوان ورشکسته طبقه بندی می شود. جدول شماره ۶ توانایی و دقت الگوهای ژنتیک فرج زاده و ژنتیک مک کی با توجه به اطلاعات شرکت های نمونه ارائه کرده است. در ستون ناموفق مدل فرج زاده عدد ۴ و درصد ۱۳/۳ نشان دهنده خطای نوع دوم و در ستون موفق نیز عدد ۱ و درصد ۳/۳ نشان دهنده خطای نوع اول است. در ستون ناموفق مدل مک کی عدد صفر و درصد صفر نشان دهنده خطای نوع دوم و در ستون موفق نیز عدد ۵ و درصد ۱۶/۷ نشان دهنده خطای نوع اول است. به این ترتیب مدل های ژنتیک فرج زاده و ژنتیک مک کی به ترتیب دارای دقت کلی ۹۱/۷ و ۹۱/۶۵ درصد می باشند.

نتیجه گیری

خلاصه نتایج حاصل از آزمون مدل ها در جدول ۷ نشان داده شده است. فرضیه اول بیان می دارد که مدل های آلتمن، شیراتا، اهلسون، زمیسکی، اسپرینگیت، سی ای اسکور، فولمر، ژنتیک فرج زاده، ژنتیک مک کی دارای قدرت پیش بینی قابل قبول توقف فعالیت در شرکت های پذیرفته شده در بورس اوراق بهادار تهران هستند. با توجه به نتایج تحقیق فرضیه اول در مورد مدل های زمیسکی، اسپرینگیت، سی ای اسکور، ژنتیک فرج زاده، ژنتیک مک کی در سطح اطمینان ۹۰ درصد پذیرفته شد.

جدول ۷- خلاصه نتایج آزمون

مدل	مدل تعدیل شده	متغیر موثر	نتیجه
مدل آلتمن	-	-	-
مدل شیراتا	-	-	-
مدل اهلسون	-	-	-
مدل فولمر	-	-	-
مدل زمیسکی	$Z=12.337B$	نسبت کل بدهی ها به کل دارایی ها	٪۹۰
مدل اسپرینگیت	$Z=8.503A$	نسبت سرمایه در گردش به کل دارایی ها	٪۹۰
مدل سی ای اسکور	$Z=12.923A$	حقوق صاحبان سهام (۱) به کل دارایی (۱)	٪۹۰
مدل ژنتیک فرج زاده	-	-	٪۹۱/۷
مدل ژنتیک مک کی	-	-	٪۹۱/۶۵

جدول ۸- مقایسه تکنیک مدل سازی و توان پیش بینی الگوها

مدل	تکنیک مدل سازی	تعدیل شده / اصلی	توان پیش بینی		
			نا موفق	موفق	کل
مدل زمیسکی	کلاسیک (پرویت)	تعدیل شده	٪۹۰	٪۹۰	٪۹۰
مدل اسپرینگیت	کلاسیک (تحلیل تشخیصی چند گانه)	تعدیل شده	٪۹۰	٪۹۰	٪۹۰
مدل سی ای اسکور	کلاسیک (تحلیل تشخیصی چند گانه)	تعدیل شده	٪۹۰	٪۹۰	٪۹۰
مدل فرج زاده	هوش مصنوعی (ژنتیک)	اصلی	٪۹۶/۷	٪۸۶/۷	٪۹۱/۷
مدل مک کی	هوش مصنوعی (ژنتیک)	اصلی	٪۸۳/۳	٪۱۰۰	٪۹۱/۶۵

با توجه به نتایج به دست آمده از آزمون مدل ها فرضیه دوم نیز مورد پذیرش قرار گرفت. به این ترتیب که مدل هایی که با استفاده از تکنیک های هوش مصنوعی (الگوریتم ژنتیک) مدل

سازی شده اند نسبت به مدل هایی که با استفاده از تکنیک های آماری مدل سازی شده اند (مدل های کلاسیک)، در پیش بینی ورشکستگی از قابلیت بیشتری برخوردار می باشند. همچنین با نگاهی اجمالی به جدول ۸ و مقایسه نتایج دو مدل ژنتیک آزمون شده درمی یابیم که با وجود آنکه توان پیش بینی کلی مدل فرج زاده و مک کی تقریباً برابر می باشد، اما در مدل فرج زاده خطای نوع اول، که برای ما از اهمیت بیشتری برخوردار است، به مراتب کمتر است. به عبارت دیگر این مدل توان بالایی در پیش بینی شرکت های ناموفق دارد و از میان ۳۰ شرکت ناموفق نمونه آزمایشی ۲۹ مورد را به صورت صحیح طبقه بندی نموده است. به صورت کلی نتایج این تحقیق نشان داد که پیش بینی پدیده ورشکستگی در محیط اقتصادی ایران امکان پذیر است. همچنین، از آنجایی که این پیش بینی بر اساس اطلاعات مالی موجود در صورت های مالی شرکتها انجام گرفته است، خود می تواند دلیلی بر وجود محتوای اطلاعاتی صورت های مالی برای کاربرد بهینه تر بازار سرمایه باشد.

پیشنهاد های تحقیق

پیشنهاد هایی که از نتایج این تحقیق حاصل می شود به این شرح می باشد: به سرمایه گذاران، بانکها، دولت، حسابرسان و سایر استفاده کنندگان اطلاعات حسابداری توصیه می شود برای ارزیابی شرکتهای بورس اوراق بهادار تهران و تصمیم گیری در رابطه با خرید سهام این شرکتهای، اعطای وام به این شرکتهای، ارزیابی عملکرد و اعلام تدوام فعالیت این شرکتهای از مدل های ژنتیک آزمون شده در این تحقیق استفاده کنند و جهت تحقیقات آتی پیشنهاد می گردد: از تکنیک های ناپارامتریک و متعلق به گروه هوش مصنوعی جهت ایجاد مدل های پیش بینی ورشکستگی استفاده و نتایج حاصله از مدل ایجاد شده با مدل های ژنتیک آزمون شده در این تحقیق مقایسه شود.

منابع داخلی

۱. اعتمادی، حسین، (۱۳۸۷)، "مروری بر مدل های پیش بینی ورشکستگی"، نشریه حسابدار، شماره ۵، صفحه ۳۹
۲. دلاور، علی، (۱۳۷۶)، مبانی نظری و عملی پژوهش در علوم انسانی و اجتماعی، تهران، انتشارات رشد، چاپ دوم
۳. خاکی، غلامرضا، (۱۳۸۶)، روش تحقیق با رویکردی به پایان نامه نویسی، تهران، انتشارات بازتاب، چاپ سوم
۴. فرج زاده دهکردی، حسن، (۱۳۸۶)، "کاربرد الگوریتم ژنتیک در مدل بندی پیش بینی ورشکستگی"، پایان نامه کارشناسی ارشد، دانشگاه تربیت مدرس، دانشکده علوم انسانی
۵. کمیته تدوین استانداردهای حسابداری، (۱۳۸۶)، استانداردهای حسابداری، تهران، سازمان حسابرسی، چاپ ششم

منابع خارجی

6. Altman, E. I, (1993), **Corporate Financial Distress and Bankruptcy: A Complete Guide to predicting and Avoiding Distress and Profiting from Bankruptcy**, second edition, John Wiley and Sons.

7. Beaver, W.H., McNichols, M.F., Rhie, J.W., (2005). "Have Financial Statements Become Less Informative? Evidence from the Ability of Financial Ratios to Predict Bankruptcy", **Review of Accounting Studies**, 10, 93–122.
8. Dimitras, A.I., Zanakis, S.H., Zopounidis, C., (1996). "A survey of business failure with an emphasis on prediction methods and industrial application", **European Journal of Operational Research**, 90, 487–513.
9. Fulmer, John G. Jr., Moon, James E., Gavin, Thomas A., Erwin, Michael J., (1984), "A Bankruptcy Classification Model For Small Firms". **Journal of Commercial Bank Lending**, pp: 25-37.
10. Lens berg T., Eilifsen, A., and McKee, T.E., (2006), "Bankruptcy theory development and classification via genetic programming", **European Journal of Operational Research**, 169, pp.677–696.
11. Zmijewski, M. E. (1984). "Methodological Issues Related to the Estimation of Financial Distress Prediction Models", **Journal of Accounting Research**, Vol: 24 (Supplement), pp: 59-82.

-
- 2-Linear probability model
 - 3-Logit
 - 4-Probit
 - 5-Partial adjustment processes
 - 6-Artificial intelligence techniques
 - 7-Recursive partitioning algorithm (decision trees)
 - 8-Case-based reasoning
 - 9-Artificial neural network
 - 10-Rough sets model
 - 11-Genetic algorithms
 - 12-Teoretical models