

فصلنامه پژوهشها و سیاستهای اقتصادی

سال هجدهم، شماره ۵۷، بهار ۱۳۹۰، صفحات ۱۹۸ - ۱۷۱

## پیش بینی تقاضای تجهیزات پزشکی (سی تی اسکن) بر اساس شبکه های عصبی مصنوعی و روش ARIMA

احمد جعفر نژاد

استاد دانشکده مدیریت دانشگاه تهران

jafarnjd@ut.ac.ir

محسن سلیمانی

کارشناس ارشد مدیریت صنعتی

rostan\_ms@yahoo.com

بخش بهداشت و درمان و زیرساخت های مورد نیاز آن هم در بخش نرم افزاری و هم در بخش سخت افزاری همواره مورد توجه بوده است. در این میان اهمیت تجهیزات و اقلام پزشکی در سیستم سلامت کشور بر هیچ کس پوشیده نیست. سازمان ها و شرکت های فعال در این بخش باید بتوانند تصمیمات صحیح را با توجه به اطلاعات موجود در محیط پرنوسان کسب و کار امروز اخذ نمایند. بنابراین، تخمین مقدار تقاضا در دوره های آتی موضوعی حیاتی به نظر می رسد. روش و ابزارهای مختلفی برای انجام پیش بینی تقاضا وجود دارد که هر یک مزیت ها و نقاط ضعف مخصوص به خود را دارند. در این مقاله با استفاده از یک شبکه عصبی مصنوعی چند لایه پیشخور با دو لایه پنهان که با الگوریتم ژنتیک به عنوان الگوریتم یادگیری آموزش داده شده است، سیستمی مقایسه ای با روش رایج مورد استفاده در پیش بینی (روش باکس - جنکینز) با مدل  $ARIMA(2,1,1)$  برای پیش بینی تقاضای دستگاه سی تی اسکن ارائه شده است که با توجه به معیار سنجش دقت مدل ها یعنی میانگین مجذور خطا (MSE)، مدل شبکه عصبی اثربخشی و کارایی بیشتری را در مقابل با روش آریمای در پیش بینی تقاضای دستگاه سی تی اسکن با توجه به داده ها و اطلاعات موجود از خود نشان داده است.

طبقه بندی JEL: C4, Q31.

واژه های کلیدی: پیش بینی تقاضا، مدل آریمای، شبکه های عصبی مصنوعی، الگوریتم ژنتیک، تجهیزات پزشکی.

تاریخ پذیرش: ۱۳۹۰/۱/۲۳

\* تاریخ دریافت: ۱۳۸۹/۸/۱۷

## ۱. مقدمه

امروزه سازمان‌ها و شرکت‌های فعال در کسب و کار با یک محیط رقابتی و دائماً در حال نوسان نسبت به گذشته مواجه هستند (پورتر و استرن، ۲۰۰۱ و الرام، ۱۹۹۱). این عامل افراد را با موقعیت‌های تصمیم‌گیری زیادی در این محیط رو به رو می‌کند.

یکی از این موقعیت‌های تصمیم‌گیری، نیازسنجی یا پیش‌بینی تقاضا می‌باشد که پیش‌بینی، فرایند برآورد موقعیت‌های ناشناخته است (آرمسترانگ، ۲۰۰۱). برای موفقیت در دنیای متغیر امروز، تصمیم‌های سازمان‌های فعال در کسب و کار متکی به پیش‌بینی‌های انجام شده با حداقل خطا می‌باشد که این در گرو داشتن یک سیستم پیش‌بینی مناسب است (آبراهام و لدالتر، ۱۹۸۳).

بخش بهداشت و درمان و زیرساخت‌های مورد نیاز آن هم در بخش نرم‌افزاری و هم در بخش سخت‌افزاری همواره مورد توجه مسئولان ذیربط بوده است. در این میان، اهمیت تجهیزات و اقلام پزشکی در سیستم سلامت کشور بر هیچ‌کس پوشیده نیست. با پیشرفت تکنولوژی و ساخت تجهیزات پزشکی جدید، سرمایه و هزینه‌های قابل‌توجهی در این بخش صرف می‌شود. کشور ما دارای رشد جمعیتی بین ۳/۱ تا ۵/۱ درصد در سال می‌باشد. با توجه به سیاست‌های دولت مبتنی بر ارتقاء سیستم بهداشت و درمان، رشد شهرنشینی، صنعتی‌شدن، تقاضای بالا برای درمان و افزایش توجه به سلامت از سوی مردم، توسعه بیمارستان‌ها و مراکز بهداشتی درمانی در سرتاسر کشور و تجهیز و نوسازی امکانات سخت‌افزاری پزشکی و درمانی آنها از یک سو و تقاضا برای درمان از سوی بیماران از کشورهای دیگر به دلیل هزینه‌های پایین درمانی در ایران نسبت به سایر کشورهای منطقه، بازار ۵۶۰ میلیون دلاری و افزایش توجه به تحقیقات در زمینه‌های مختلف علوم پزشکی در راستای اهداف سند چشم‌انداز بیست ساله از سوی دیگر، باعث شده است تا تقاضا برای پیشگیری، تشخیص و درمان از سوی مردم افزایش یابد که این امر افزایش تقاضا برای امکانات سخت‌افزاری و نرم‌افزاری را در این بخش در پی داشته است.

امروزه شبکه‌های عصبی به عنوان یکی از مؤثرترین روش‌های هوش محاسباتی در شاخه‌های مختلف علوم شناخته شده‌اند. این شبکه‌ها الهام گرفته از نحوه کارکرد سیستم عصبی انسان می‌باشد. در واقع، کارایی در خور توجهی که دستگاه عصبی انسان در درک و تشخیص پدیده‌های مختلف از خود نشان داده است محققان را در علوم مختلف از جمله مدیریت بر آن داشته که با تقلید از این شبکه اقدام به طراحی شبکه‌های عصبی مصنوعی برای تجزیه و تحلیل مسائل مختلف بنمایند. شبکه‌های عصبی مصنوعی به دلیل هوشمندبودن، سرعت بالای پردازش داده‌ها، قابلیت تطبیق با تغییرات محیطی، قابلیت تعلیم و قابلیت مدل‌سازی سیستم‌های غیرخطی با پیچیدگی زیاد (فاوست، ۱۹۹۴) و الگوریتم ژنتیک به

دلیل جستجوی چندجانبه، استفاده از قوانین احتمالی، طبیعت تکاملی، سرعت بالا در ارائه نتایج قابل قبول (سیوانان دام و دپا، ۲۰۰۸) با توجه به پیچیدگی‌های مسائل و محیط پیش‌روی سازمان‌ها از ابزارهای قوی، هوشمند و توانایی برای استفاده در تصمیم‌گیری‌ها می‌باشند.

به این منظور، این مقاله یک روش سیستماتیک را برای انجام پیش‌بینی تقاضا در بخش تجهیزات پزشکی برای شرکت‌های فعال در این بخش ارائه می‌نماید. بخش دوم، به بررسی تحقیقات صورت گرفته در خصوص پیش‌بینی تقاضای تجهیزات پزشکی می‌پردازد. بخش سوم، مدل‌های مورد استفاده در روش پیشنهاد شده را توصیف می‌کند. بخش چهارم، نتایج حاصل از مدل‌های مورد استفاده را ارائه می‌دهد. در بخش پایانی، پیشنهادها برای تحقیقات آتی در رابطه با این موضوع ارائه می‌شود.

## ۲. پیشینه تحقیق

در رابطه با پیش‌بینی تقاضا در بخش تجهیزات پزشکی تاکنون تحقیق ثبت شده‌ای به صورت مقاله انجام نشده است، اما در رابطه با استفاده از شبکه‌های عصبی مصنوعی در پیش‌بینی تقاضا در بخش‌های مختلف تحقیقات بسیاری انجام شده است که در ادامه به تعدادی از آنها اشاره می‌شود.

افندیگل، اونات و قهرمان (۲۰۰۹) به پیش‌بینی تقاضای زنجیره تأمین یک شرکت فروشنده کالاهای مصرفی با دوام در ترکیه با شبکه‌های عصبی مصنوعی و فازی پرداخته و نشان دادند که شبکه عصبی فازی دارای عملکرد بهتری در پیش‌بینی تقاضا است.

لوهاج، ریس و استنس بال (۱۹۹۶) از یک روش ترکیبی مدل‌سازی شبکه عصبی - اقتصاد سنجی برای پیش‌بینی فروش یک شرکت فروشنده کالاهای مصرفی در دانمارک استفاده کردند. این مدل تلاش کرد تا خصوصیات ساختاری مدل‌های اقتصادسنجی را با خصیصه‌های بازشناسی الگوی غیرخطی شبکه عصبی ادغام کند.

آبورتو و وبر (۲۰۰۷) یک سیستم هوشمند ترکیبی را با ادغام مدل‌های ARIMA و شبکه عصبی برای پیش‌بینی تقاضا در زنجیره تأمین و توسعه یک سیستم جایگزینی در یک سوپرمارکت واقع در شیلی ارائه کردند.

کائو و ژیو (۱۹۹۸) یک سیستم هوشمند پیش‌بینی تقاضا را از طریق ادغام شبکه‌های عصبی فازی و شبکه‌های عصبی مصنوعی برای استفاده در یک سوپرمارکت در چین ارائه کردند.

کائو (۲۰۰۱) یک سیستم پیش‌بینی فروش را بر اساس شبکه‌های عصبی فازی با الگوریتم یادگیری ژنتیک برای یک شرکت CVS در تایوان ارائه و نتایج را با مدل ARMA مقایسه و نشان داد که شبکه عصبی فازی دارای عملکرد بهتری است.

کائو، وو و وانگ (۲۰۰۱) یک سیستم پیش‌بینی فروش را بر اساس ادغام شبکه‌های عصبی مصنوعی و شبکه‌های عصبی فازی با حذف وزن‌های فازی شرکت CVS تایوانی ارائه و مانند مدل قبلی نتایج را با مدل ARMA مقایسه و نشان دادند که مدل ترکیبی دو شبکه دارای عملکرد بهتری می‌باشد. تنگ، آلمیدا و فیش وک (۱۹۹۱) مدل‌های باکس و جنکینز و شبکه‌های عصبی مصنوعی را برای پیش‌بینی رفت‌وآمد مسافر خط هوایی بین‌المللی، پیش‌بینی فروش ماشین در داخل آمریکا و پیش‌بینی فروش ماشین در خارج از آمریکا بکار بردند و نشان دادند که مدل‌های باکس و جنکینز در کوتاه‌مدت بهتر از شبکه‌های عصبی مصنوعی هستند و در بلندمدت شبکه‌های عصبی مصنوعی بهتر هستند.

سان، چوئی و یونگ یو (۲۰۰۸) به بررسی پیش‌بینی فروش به وسیله شبکه‌های عصبی مصنوعی در یک خرده‌فروشی فعال در عرصه مد در هنگ‌کنگ پرداختند.

چکراپورتی، مهروترا و موهان (۱۹۹۲) به بررسی پیش‌بینی قیمت آرد در سه شهر آمریکا با شبکه‌های عصبی مصنوعی برای تجزیه و تحلیل سری زمانی چند متغیره پرداختند و نشان دادند که روش پیشنهادی دارای عملکرد بهتری نسبت به روش‌های سنتی می‌باشد.

آلون، کیو آی و سادوفسکی (۲۰۰۱) به مقایسه پیش‌بینی فروش خرده‌فروشی‌های تجمیعی در آمریکا با شبکه‌های عصبی مصنوعی و روش‌های رایج سنتی پرداخته و نشان دادند که مدل شبکه‌های عصبی مصنوعی دارای عملکرد بهتری نسبت به روش‌های سنتی می‌باشد.

لچترماچر و فولر (۱۹۹۵) یک مدل شبکه عصبی مصنوعی کالیره را ایجاد کردند. مدل روش‌های باکس - جنکینز را استفاده می‌کردند تا با شناسایی وقفه‌های داده‌ها آنها را به عنوان متغیر ورودی استفاده نمایند. آنها از روش کاوشی برای شناسایی واحدهای پنهان مورد نیاز در ساختار شبکه استفاده می‌کردند.

دوو و ولف (۱۹۹۷) جزئیات پیاده‌سازی شبکه‌های عصبی و یا سیستم‌های منطق فازی را در صنعت به ویژه در بخش‌های زمانبندی و برنامه‌ریزی، کنترل موجودی، کنترل کیفیت، فناوری گروهی و پیش‌بینی ارائه کردند.

از دیگر موارد می‌توان به پیش‌بینی تقاضای کوتاه‌مدت و بلندمدت جریان برق (ال - سبا و ال - امین، ۱۹۹۹، بکالی، سلورا، برانو و مارواگلیا، ۲۰۰۴)، پیش‌بینی تقاضای استفاده از انرژی (هابس، جیتراپایکولسارن، کوندا و ماراچیکولام، ۱۹۹۸، سوزن، آرککلیوگلا و ازکایماک، ۲۰۰۵)، پیش‌بینی تقاضای تورسیم (لاو، ۲۰۰۰، لاو و آئو، ۱۹۹۹ و پالمر، مونتانو و سس، ۲۰۰۶)، پیش‌بینی ورشکستگی

ریسک اعتباری (آتیا، ۲۰۰۱)، پیش‌بینی نرخ مبادلات ارزهای خارجی (هانگ، لای، ناکاموری و وانگ، ۲۰۰۴) و ... اشاره کرد. در تمام این موارد شبکه‌های عصبی بکار رفته عملکرد بهتری نسبت به روش‌ها و مدل‌های سنتی داشته‌اند.

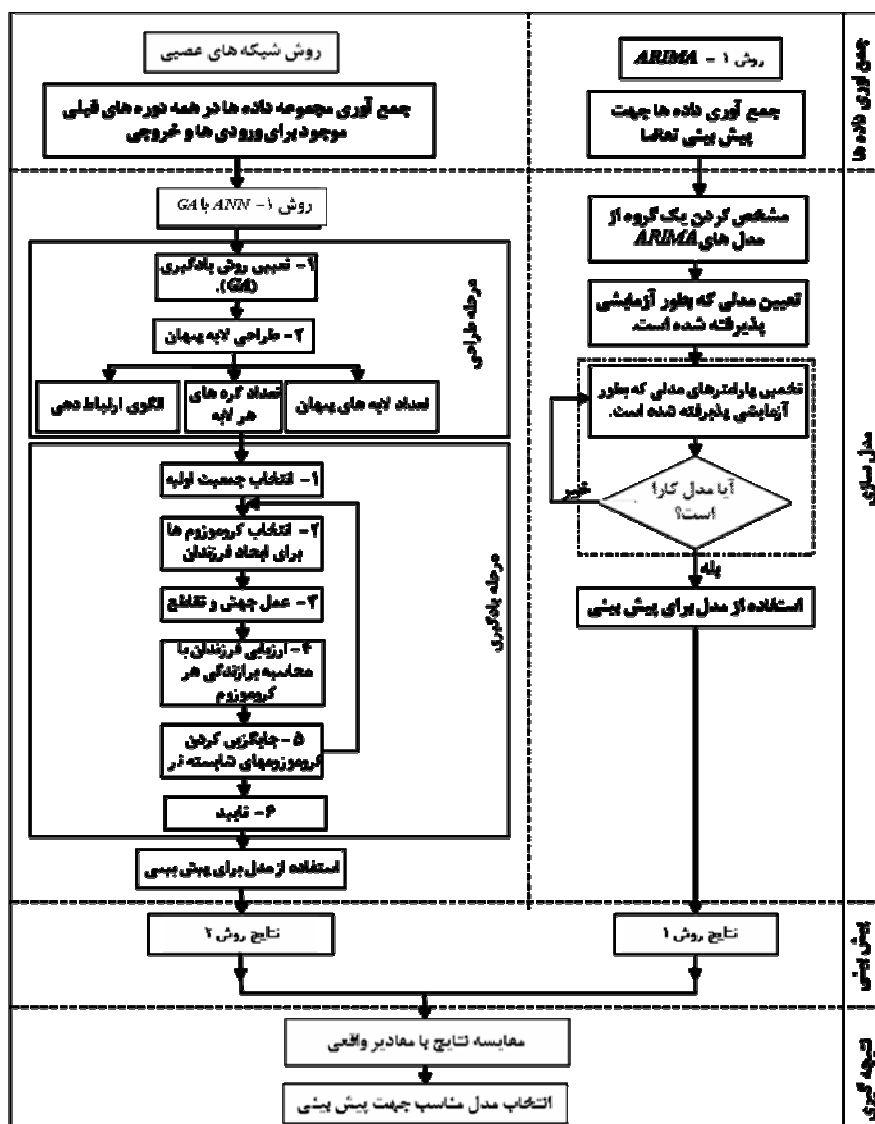
### ۳. روش تحقیق

فنون پیش‌بینی هوش مصنوعی اخیراً به منظور حل مسائل مدیریتی مورد توجه بسیاری از محققان قرار گرفته است. آنها این توانایی را دارند تا همانند انسان‌ها از طریق اندوختن دانش از طریق فعالیت‌های یادگیری تکراری یاد بگیرند. بنابراین، هدف این مقاله ارائه سیستمی برای پیش‌بینی تقاضای تجهیزات پزشکی برای مدیریت تقاضا در محیط پر نوسان امروز می‌باشد. به این منظور، در این مقاله یک سیستم مقایسه‌ای بر اساس شبکه‌های عصبی چندلایه پیشخور که با الگوریتم ژنتیک آموزش داده شده است و روش ARIMA برای این کار ارائه شده است. نمودار (۱) فرایند مورد استفاده در پیش‌بینی را نشان می‌دهد. در ادامه، روش‌های مورد استفاده تشریح می‌شوند.

#### ۳-۱. روش ARIMA

باکس و جنکینز<sup>۱</sup> در سال ۱۹۶۰ روشی را توسعه دادند که اساس آن بر مبنای بررسی حوزه وسیعی از مدل‌های پیش‌بینی برای یک سری زمانی قرار گرفته است. این روش مدل خود را از میان مجموعه‌ای از مدل‌ها که به آن ARIMA می‌گویند، با روشی سیستماتیک انتخاب می‌کند. در مدل ARIMA فرض می‌شود که مقادیر آتی یک متغیر تابعی خطی از چندین مشاهده گذشته و خطاهای تصادفی می‌باشند. این مدل‌ها و ابزارهای استفاده‌شده در آن تنها برای سری‌های زمانی ایستا (مانا) کاربرد دارد. بنابراین، پیش از تحلیل یک سری زمانی غیرایستا به وسیله این مدل می‌بایست با استفاده از روش‌های دیفرانسیل‌گیری به یک سری ایستا تبدیل شود.

1. Box & Jenkins



نمودار ۱. فرایند مورد استفاده برای پیش بینی تجهیزات پزشکی

صورت کلی مدل آریما به این شکل می باشد:

$$\varphi(B)(1-B)^d X_t = \theta(B)Z_t \quad (1)$$

$$\varphi(B) = 1 - \varphi_1 B - \varphi_2 B^2 - \varphi_3 B^3 - \dots - \varphi_p B^p \quad (2)$$

$$\theta(B) = 1 + \theta_1 B + \theta_2 B^2 + \theta_3 B^3 + \dots + \theta_q B^q \quad (3)$$

که در آن،  $Z_t, X_t, \phi(B), \theta(B), p, q, d, \phi_i (i = 1, 2, \dots, p)$  و  $\theta_j (j = 1, 2, \dots, q)$  به ترتیب مقادیر آتی متغیر، نویز سفید در زمان  $t$ ، چند جمله‌ای اتورگرسیو، چند جمله‌ای میانگین متحرک، مرتبه اتورگرسیو، مرتبه میانگین متحرک، درجه تفاضل‌گیری و پارامترهای مدل‌های اتورگرسیو و میانگین متحرک می‌باشند.

مدل‌های ARIMA طی سه مرحله ساخته می‌شوند. در مرحله اول تعدادی مدل از بین گروه عمومی مدل‌های ARIMA بر اساس معیارهای مربوط بدون قواعد مشخص و بر اساس قضاوت و تجربه تحلیلگر شناسایی می‌شوند که در این مقاله برای شناسایی نوع و مرتبه مدل از نمودار تابع خود همبستگی (ACF) و تابع خودهمبستگی جزئی سری (PACF) (در صورت نیاز از سری تفاضل‌گیری شده) برای ۱۶ وقفه (Lag) و برای سنجش ایستایی و حذف روند از آزمون تعمیم‌یافته دیکی - فولر (دیکی و فولر، ۱۹۷۹) استفاده می‌شود.

مرحله دوم به تخمین و آزمون پارامترهای مدل‌های شناسایی شده می‌پردازد که در اینجا از روش حداقل مربعات خطی استفاده شده و معنادار بودن ضرایب در سطح اطمینان ۹۵ درصد آزمون می‌شود. برای انتخاب مناسب‌ترین مدل از معیار اطلاعات آکائیک<sup>۱</sup> (۱۹۷۹)، معیار اطلاعات شوارتز<sup>۲</sup> (۱۹۷۸)، خطای پیش‌بینی نهایی آکائیک<sup>۳</sup> (۱۹۷۰)، مجموع مربعات خطا<sup>۴</sup>، میانگین قدر مطلق درصد خطا<sup>۵</sup> و ریشه میانگین مربعات خطا<sup>۶</sup> استفاده شده است. با توجه به تعدد معیارها برای انتخاب بهترین مدل از روش TOPSIS استفاده شده است.

مرحله سوم، مرحله بازبینی تشخیص مدل (آزمون پسماندها) می‌باشد. در این مرحله کارایی (بررسی تصادفی بودن خطاهای مدل مورد بررسی) سنجیده می‌شود که در این مقاله از آزمون الجانگ - باکس<sup>۷</sup> (۱۹۷۸) استفاده شده است. در صورتی که مدل برازش شده، کارا نباشد به گروه مدل‌های ARIMA برگشته و مدل دیگر از آنها را انتخاب می‌کنیم. نهایتاً بعد از تشخیص بهترین مدل، با استفاده از آن به پیش‌بینی مقادیر آینده سری زمانی می‌پردازیم (چتفیلد، ۲۰۰۴).

1. Akaike Information Criteria (AIC)
2. Schwartz Bayesian Criterion (SBC)
3. Final Prediction Error (FPE)
4. Sum of Squares Error (SSE)
5. Mean Absolute Percentage Error (MAPE)
6. Root Mean Square Error (RMSE)
7. Ljung-Box Test

## ۳-۲. روش شبکه عصبی

شبکه عصبی مصنوعی یک سیستم پردازش اطلاعات است که ویژگی‌های عملکردی مشابه شبکه عصبی بیولوژیکی دارد (فاوست، ۱۹۹۴). طی دو دهه گذشته، شبکه‌های عصبی مصنوعی کارایی خود را در مدل‌سازی مسائل پیچیده و نامفهوم به اثبات رسانیده‌اند (دار و استین، ۱۹۹۷). شبکه عصبی مصنوعی یک فناوری می‌باشد که عمدتاً برای پیش‌بینی، خوشه‌بندی، طبقه‌بندی و هشدار دادن برای الگوهای غیرعادی استفاده شده‌اند (هاکینز، ۱۹۹۴).

شبکه‌های عصبی چندلایه پیشخور<sup>۱</sup> شبکه رایج در خصوص مسائل پیش‌بینی هستند که در این مقاله مورد استفاده قرار گرفته است. عوامل متعددی مانند تعداد لایه‌های پنهان، تعداد نرون‌های لایه‌های پنهان، نرمال کردن داده‌ها و تعیین الگوریتم یادگیری می‌توانند در عملکرد شبکه‌های عصبی مؤثر باشند. پس معماری مناسب شبکه با استفاده از تجربه، آزمایش و خطا حاصل می‌شود که در ادامه این موارد مشخص می‌شوند.

## ۳-۲-۱. معماری شبکه

## ۳-۲-۱-۱. تعیین تعداد لایه‌های پنهان

این پارامتر به پیچیدگی مسأله و تعداد متغیرها وابسته است. محققان بیشتر از یک لایه پنهان برای مسائل پیش‌بینی استفاده کرده‌اند (زانگ، پاتوو و هو، ۱۹۹۸). استفاده از دو لایه پنهان می‌تواند موجب بهبود کارایی در فرایند یادگیری شبکه شود (اسرینیواسون، چانگ و لیو، ۱۹۹۴). در این مقاله مقدار این متغیر یک بار یک و یک بار دو در نظر شده است.

## ۳-۲-۲. تعیین تعداد نرون‌های لایه‌های پنهان

تعیین این پارامتر کاری پیچیده و بدون مبنای تئوریک است. در کل، شبکه با نرون‌های کمتر در لایه‌های پنهان معمولاً توانایی تعمیم‌پذیری بهتری دارد و کمتر دچار برازش بیش از حد می‌شود (زانگ و دیگران، ۱۹۹۸). با افزایش تعداد نرون‌های لایه پنهان، توان شبکه در تشخیص پیچیدگی‌های موجود در مجموعه آموزشی افزایش می‌یابد اما قابلیت تعمیم شبکه کاهش می‌یابد. چندین قانون سرانگشتی برای این منظور پیشنهاد شده است. در این مقاله برای شبکه عصبی برای لایه‌های پنهان از قانون  $n$  (تنگ و فیش وک، ۱۹۹۳) (که  $n$  نشان دهنده تعداد نرون‌های لایه ورودی است) استفاده شده است.

---

1. Multi Layer Perceptron

## ۳-۲-۱. تعیین تعداد نرون‌های لایه خروجی

تعیین این پارامتر نسبتاً آسان می‌باشد و مستقیماً به مسأله مورد مطالعه مربوط می‌شود. در مسائل پیش‌بینی این پارامتر اغلب با افق پیش‌بینی تطابق دارد. در مسائلی که پیش‌بینی یک دوره بعدی مطرح است، تعداد نرون‌های لایه خروجی را یک در نظر می‌گیرند. در اینجا چون یک خروجی (میزان تقاضا) وجود دارد، بنابراین تعداد نرون‌های لایه خروجی یک در نظر گرفته شده است.

## ۳-۲-۴. تعیین توابع فعال‌سازی نرون‌ها

تابع فعال‌سازی (تابع انتقال) رابطه بین ورودی‌ها و خروجی یک نرون و یک شبکه را تعیین می‌کند. یک شبکه ممکن است از توابع فعال‌سازی متفاوت برای نرون‌های متفاوت در لایه‌های متفاوت یا مشابه استفاده کند (شانبرگ، ۱۹۹۰). محققان بیشتر از توابع انتقال زیگموئیدی<sup>۱</sup> برای نرون‌های لایه‌های پنهان استفاده می‌کنند که رایج‌ترین تابع انتقال می‌باشد (کائو، ۲۰۰۱). در این مقاله برای تمام نرون‌ها در یک لایه از تابع انتقال مشابه استفاده شده است. تابع فعال‌سازی نرون‌های لایه‌های پنهان "تانزانت هایربولیک زیگموئیدی" می‌باشد و تابع فعال‌سازی نرون لایه خروجی خطی می‌باشد.

## ۳-۲-۵. نرمال کردن داده‌ها

نرمال‌ایز کردن داده‌ها برای اینکه داده‌های بزرگتر داده‌های کوچکتر را تحت شعاع قرار ندهند و همچنین برای جلوگیری از اشباع زود هنگام نرون‌های لایه‌های پنهان که مانع یادگیری شبکه عصبی می‌شود ضروری است (هاجر و بشر، ۲۰۰۰). هیچ رویه استاندارد برای نرمال کردن ورودی‌ها و خروجی‌ها وجود ندارد. در این مقاله از نرمال‌سازی ساده برای نرمال کردن داده‌ها استفاده شده است.

$$X_n = \frac{X}{X_{\max}} \quad (5)$$

که در آن،  $X_n$ ،  $X$  و  $X_{\max}$  به ترتیب ارزش نرمال‌شده، ارزش واقعی متغیر ورودی و حداکثر ارزش متغیر ورودی می‌باشد.

## ۳-۲-۶. تعیین داده‌های مربوط به آموزش و تست

مجموعه داده‌های آموزش برای توسعه مدل و مجموعه داده‌های تست<sup>۲</sup> برای ارزیابی توانایی پیش‌بینی مدل انتخاب می‌شود (وی‌جند، هابرمین و راملهارت، ۱۹۹۲). هیچ راه حل کلی برای تعیین این دو

1. Sigmoid Transfer Function  
2. Test Sample

مجموعه وجود ندارد، اما چندین تجربی روش برای تعیین این مجموعه‌ها پیشنهاد شده است. در این مقاله از روش (۸۰/۲۰) برای تعیین داده‌های آموزش و تست استفاده شده است. یعنی ۸۰ درصد داده‌ها مربوط به داده‌های آموزش و ۲۰ درصد مربوط به داده‌های تست می‌باشد.

### ۳-۲-۱. تعیین معیار عملکرد شبکه

مهم‌ترین معیار عملکرد دقت پیش‌بینی است که به وسیله آن می‌توان به ماورای داده‌های آموزش دسترسی پیدا کرد. معیار رایج برای سنجش عملکرد شبکه در مسائل پیش‌بینی MSE<sup>۱</sup> می‌باشد که در این مقاله استفاده شده است.

$$MSE = \frac{1}{n} \sum_{t=1}^n (Y_t - \hat{Y}_t)^2 \quad (۶)$$

### ۳-۲-۲. طراحی الگوریتم ژنتیک

مهم‌ترین موضوعی که شبکه‌های عصبی را از دیگر روش‌های بهینه‌سازی جدا می‌کند، یادگیری یا آموزش دیدن آنها می‌باشد. در حالت کلی دو نوع یادگیری وجود دارد یکی یادگیری با ناظر<sup>۲</sup> و دیگری یادگیری بدون ناظر<sup>۳</sup>.

رایج‌ترین الگوریتم، یادگیری برای شبکه‌های چندلایه پیشخور، الگوریتم پس انتشار خطا<sup>۴</sup> می‌باشد که یک روش یادگیری با نظارت است. در این مقاله، از الگوریتم ژنتیک برای مرحله یادگیری استفاده شده است. الگوریتم ژنتیک یک روش جستجو بر حل مسائل بهینه‌سازی است که بر اساس مفهوم تکامل می‌باشد (هالند، ۱۹۷۵) و در زمره روش‌های یادگیری بدون ناظر است. مراحل انجام الگوریتم ژنتیک برای یادگیری در شبکه به قرار زیر است:

- پارامترهای جمعیت: مورد اول مربوط به ایجاد جمعیت اولیه می‌شود. در این مقاله، جمعیت اولیه به صورت تصادفی با یک توزیع یکنواخت ایجاد شده است. مورد بعدی مربوط به اندازه جمعیت می‌شود و این پارامتر تعداد افراد را در جمعیت بیان می‌کند. معمولاً مقدار آن را بین ۵۰۰ - ۲۰ در نظر می‌گیرند. در این مقاله، مقدار این پارامتر ۲۰، ۵۰، ۸۰، ۱۰۰، ۱۲۰، ۱۵۰ و ۱۸۰ برای یافتن مقدار بهینه در نظر گرفته شده است. پارامتر بعدی تعداد نسل‌ها<sup>۵</sup> می‌باشد. این پارامتر تعداد تکرارهای الگوریتم

1. Mean Squared Error
2. Supervised Learning
3. Unsupervised Learning
4. Back Propagation
5. Number of Generation

ژنتیک را نشان می‌دهد و یکی از پارامترهای توقف اجرای الگوریتم می‌باشد. در این مقاله، مقدار این پارامتر ۲۰۰ در نظر گرفته شده است.

- تابع برازش: برای معرفی تابع برازش می‌بایست متغیرها در مدل قرار داده شوند و سپس تفاوت بین مقادیر تخمین‌زده شده و داده‌های واقعی برای هر فرد محاسبه شود. در هر نسل (تکرار) افراد با حداقل تفاوت باید بازگردانده شود. در این تحقیق، تابع برازش به صورت زیر تعریف می‌شود:

$$\text{Fitness Function} = 1 / \left( \frac{1}{n} \sum_{t=1}^n (Y_t - \hat{Y}_t)^2 \right) \quad (7)$$

که در آن،  $Y_t$ ،  $\hat{Y}_t$  و  $n$  به ترتیب مقدار واقعی، مقدار پیش‌بینی و تعداد مشاهدات می‌باشد.

- انتخاب: این اپراتور یکی از سه اپراتور اصلی الگوریتم ژنتیک است. این اپراتور وظیفه انتخاب والدین را برای نسل بعد دارد. یک فرد در صورتی که بتواند ژن‌هایش را به بیش از یک فرزند بدهد می‌تواند به عنوان والد انتخاب شود. در این تحقیق از روش انتخاب مسابقه<sup>۱</sup> (گلدبرگ، ۱۹۸۹) استفاده شده است. این روش که از سری روش‌های نمونه‌گیری مختلط می‌باشد و شامل مراحل زیر است:

اندازه مسابقه را انتخاب کنید ( $k \geq 2$ ) سپس ترتیبی تصادفی از افراد موجود در جمعیت ایجاد می‌شود و مقدار برازش اولین  $k$  رشته فهرست‌شده در ترتیب با هم مقایسه شده و بهترین رشته به نسل بعد منتقل می‌شود. سپس، رشته‌های مقایسه‌شده حذف می‌شود. اگر ترتیب خالی شده باشد ترتیب دیگری ایجاد می‌شود و در نهایت مراحل ۳ و ۴ تا زمانی که برای نسل بعد به انتخاب نیاز نباشد تکرار می‌شود و معمولاً از  $k = 2$  استفاده می‌شود.

- تقاطع: این اپراتور نیز جزء اپراتورهای اصلی الگوریتم ژنتیک است. اولین پارامتر در اینجا تعیین احتمال تقاطع<sup>۲</sup> می‌باشد. عملیات تقاطع روش اصلی برای ایجاد افراد است. بنابراین، احتمال تقاطع باید نسبتاً مقدار بزرگی را اختیار کند. مقدار آن معمولاً بین  $[0.7 - 0.9]$  است. در این مقاله، مقدار آن  $P_c = 0.88$  در نظر گرفته شده است.

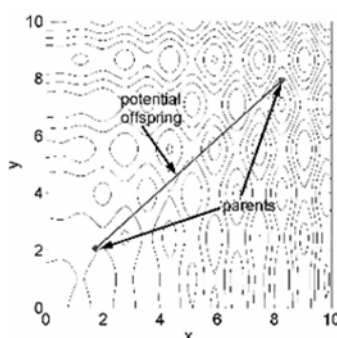
پس از این مورد، تعیین روش تقاطع قرار دارد. در این تحقیق از میان انواع روش‌های موجود برای این کار، روش تقاطع کاوشی<sup>۳</sup> انتخاب شده است. این روش از مقدار برازش دو والد برای تقاطع استفاده می‌کند. فرزندان روی خطی که بین دو والد قرار دارد به صورت زیر بوجود می‌آیند:

1. Tournament Selection
2. Crossover Probability
3. Heuristic Crossover

$$\text{Offspring1} = \text{BestParent} + r * (\text{BestParent} - \text{WorstParent}) \quad (۸)$$

$$\text{Offspring2} = \text{BestParent} \quad (۹)$$

که در آن،  $r$  یک عدد بین  $[۰ - ۱]$  می باشد. در این مقاله، مقدار آن  $۰/۹۹$  در نظر گرفته است. نمودار (۲) این موضوع را نشان می دهد.



نمودار ۲. تقاطع ابتکاری

- جهش: این اپراتور نیز جزء اپراتورهای اصلی الگوریتم ژنتیک است. اولین پارامتر در اینجا تعیین احتمال جهش<sup>۱</sup> می باشد. مقدار آن معمولاً بین  $[۰/۰۰۱ - ۰/۰۱]$  می باشد. در این تحقیق، مقدار آن  $P_m = ۰/۰۱$  در نظر گرفته است.

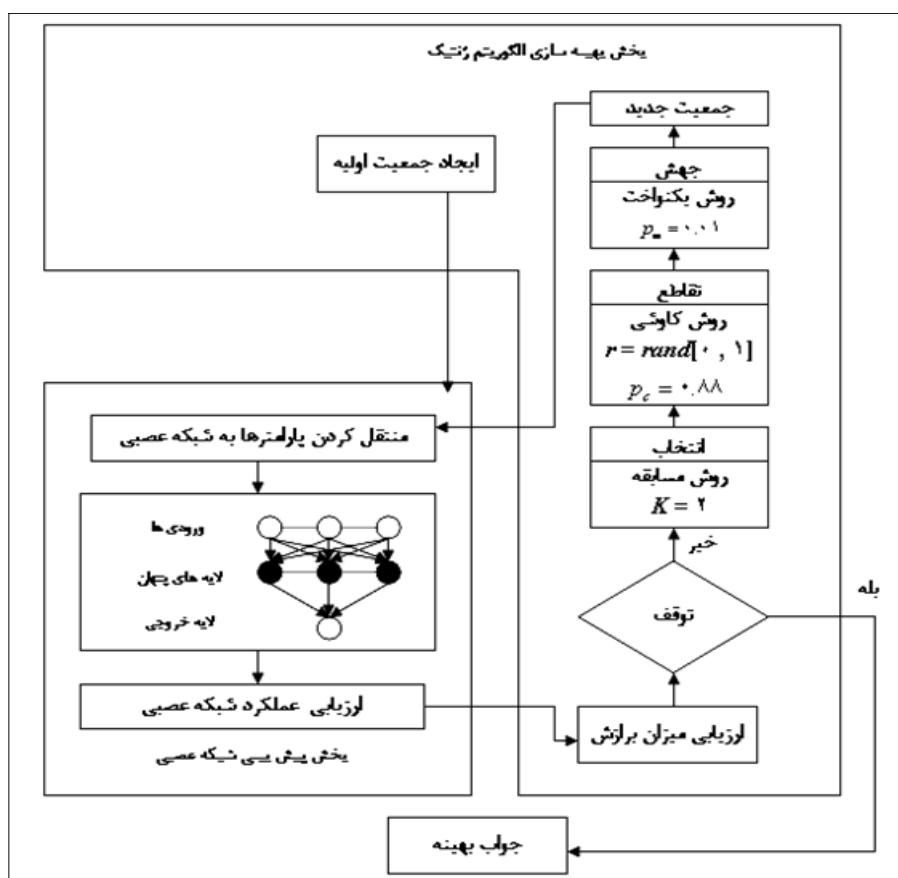
مرحله بعد، تعیین روش جهش می باشد. در این تحقیق از میان انواع روش ها موجود برای این کار، روش جهش یکنواخت<sup>۲</sup> در نظر گرفته شده است. در این روش ابتدا ژنی از کروموزوم آماده برای جهش به صورت تصادفی انتخاب شده و مقدار آن به مقدار تصادفی دیگری تغییر می یابد. یعنی یک عدد تصادفی در بازه  $[۱ و L]$  که  $L$  طول کروموزوم می باشد انتخاب شده و ژن موجود در آن مکان از کروموزوم تغییر می کند. به عنوان مثال، اگر  $R = ۵$  و  $L = ۹$  باشد داریم (نمودار ۳):

Parent : 0 1 0 0 1 0 1 1 0     $\Rightarrow$     Child : 0 1 0 0 0 0 1 1 0  
 ۱ تبدیل به ۰ می شود.

نمودار ۳. جهش یکنواخت

1. Mutation Probability
2. Uniform Mutation

نمودار (۴) چگونگی انجام یادگیری روی شبکه را نشان می‌دهد.



نمودار ۴. چگونگی انجام یادگیری روی شبکه، توسط الگوریتم ژنتیک

### ۳-۳. شاخص‌ها، داده‌ها و چگونگی تلخیص داده‌ها

تجهیزات پزشکی طیف وسیعی را در بر می‌گیرند. این طیف وسیع فعالیت بیش از ۲۵۰ شرکت را در امر واردات و صادرات در پی داشته است. البته در مواردی که تجهیزات از فناوری بالایی برخوردار است این فعالیت انحصاری می‌شود. در این تحقیق از میان تمام این تجهیزات، دستگاه‌های تصویربرداری سی‌تی‌اسکن برای سنجش تقاضای آن انتخاب شده است.

- الف) شاخص‌ها: سیستم پیش‌بینی تقاضای ارائه‌شده در این مقاله از متغیرهای زیر به عنوان ورودی‌ها و خروجی شبکه عصبی استفاده می‌کند:
- چرخه عمر فناوری (ورودی): در اینجا میزان اهمیت تغییرات در تکنولوژی مورد استفاده در تجهیزات مورد تقاضا از منظر متقاضی‌دهنده مورد توجه است و این متغیر کیفی می‌باشد.
  - کیفیت (ورودی): می‌توان گفت که کیفیت برابر با تمام علائم و ویژگی‌های محصول یا خدمت است که مربوط به توانایی ارضاء نیازهای تعیین‌شده می‌باشد و این متغیر کیفی است.
  - عمر مفید (ورودی): بیانگر مدت زمانی است که دستگاه مستهلک شود یعنی استهلاک انباشته دستگاه برابر با ارزش دفتری آن می‌شود و این متغیر کمی می‌باشد.
  - بهای تمام شده مالکیت (ورودی): شامل کتمام هزینه‌هایی می‌شود که متقاضی تجهیزات از ابتدا تحصیل دستگاه تا پایان زمان استفاده از آن می‌پردازد و این متغیر کمی است.
  - تعداد استفاده‌کنندگان از تجهیزات (ورودی): بیان‌کننده تعداد افراد استفاده‌کننده از تجهیزات مورد نظر طی یک دوره زمانی معین (ماهانه) است و این متغیر کمی می‌باشد.
  - زمان انتظار یا تأخیر برای دریافت تجهیزات (ورودی): مدت زمانی که از هنگام سفارش خرید تا هنگام دریافت تجهیزات مورد نظر طول می‌کشد و این متغیر کمی می‌باشد.
  - خدمات پس از فروش (ورودی): عبارت است از مجموعه فعالیت‌هایی که از سوی شرکت طراحی می‌شود تا سطح رضایت مشتری ارتقا یابد، از جمله آموزش، نصب و راه‌اندازی، تعمیرات و .... که میزان رضایت خدمات‌گیرنده را می‌سنجد و این متغیر کیفی می‌باشد.
  - قوانین و مقررات (ورودی): میزان اهمیت مجموعه دستورالعمل‌هایی که از جانب ارگان‌های داخلی و بین‌المللی مربوط به تهیه و پخش تجهیزات پزشکی در کشورهای مبداء و مقصد روی تقاضا دارد و این متغیر کیفی است.
  - وقفه‌های متغیر وابسته (ورودی): در اینجا برای در نظر گرفتن عامل زمان در پیش‌بینی خروجی از ورودی یک و دو دوره قبل متغیر وابسته به عنوان ورودی استفاده می‌شود تا بخش‌هایی از دینامیک زمانی وارد مدل‌سازی شود و این متغیر کمی می‌باشد.
  - میزان تقاضا (خروجی): برابر با میزان تقاضای دستگاه سی‌تی‌اسکن در دوره‌های موجود می‌باشد و این متغیر کمی می‌باشد.
- مدل ARIMA نیز از داده‌های تقاضا برای مدل‌سازی استفاده می‌کند. متغیرهای کیفی به وسیله مقیاس (۹ - ۱) سنجیده شده است با توجه به اینکه اعداد حاصل به صورت متوسط می‌باشد.

ب) تهیه و جمع‌آوری داده‌ها: داده‌های کمی مورد نیاز از طریق مراجعه حضوری به اداره گمرک، سازمان بازرگانی، اداره مرکزی رادیولوژی ایران، نقطه تجاری ایران، بیمه تأمین اجتماعی و شرکت (تأمین اجتماعی) در فاصله سال‌های (۱۳۷۱-۱۳۸۷) به دست آمده است.

در خصوص متغیرهای کیفی، داده‌های موجود از ابتدای سال ۱۳۸۲ می‌باشد. برای استفاده از داده‌های سال‌های قبل از سال ۱۳۸۲ داده‌های لازم برای متغیرهای کیفی از طریق نرم‌افزار Best Fit 1.5 شبیه‌سازی شده است.

در اینجا ۲۰۴ دوره داده وجود دارد. در شبکه عصبی از این تعداد ۲۰۲ دوره داده قابل استفاده است زیرا به دلیل تعداد وقفه‌ها، دو دوره داده از کل داده‌ها کسر می‌شود. بنابراین:

$$\text{Train Data} = 0/80 \times 202 = 162 \quad (10)$$

$$\text{Test Data} = 0/20 \times 204 = 40 \quad (11)$$

در قسمت ARIMA نیز ۸۰ درصد داده‌ها برای مدل‌سازی و مابقی برای تست مدل استفاده شده‌اند.

ج) تلخیص و تحلیل داده‌ها: در بخش ARIMA برای تحلیل داده‌ها از نرم‌افزارهای EViews 5.0 استفاده شده است. در بخش شبکه‌های عصبی از نوار ابزار شبکه‌های عصبی و الگوریتم ژنتیک نرم‌افزار MATLAB R2008b 7.7 برای بهینه‌سازی استفاده شده است.

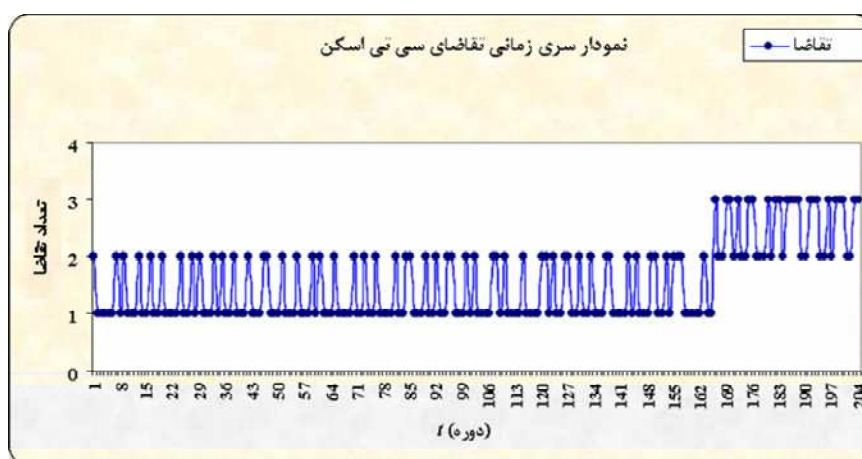
#### ۴. یافته‌های حاصل از مدل‌سازی

در این بخش ابتدا محاسبات مربوط به ساختن مدل ARIMA و سپس نتایج شبکه عصبی و تعیین شبکه بهینه ارائه می‌شود.

##### ۴-۱. محاسبات مربوط به ساختن مدل ARIMA

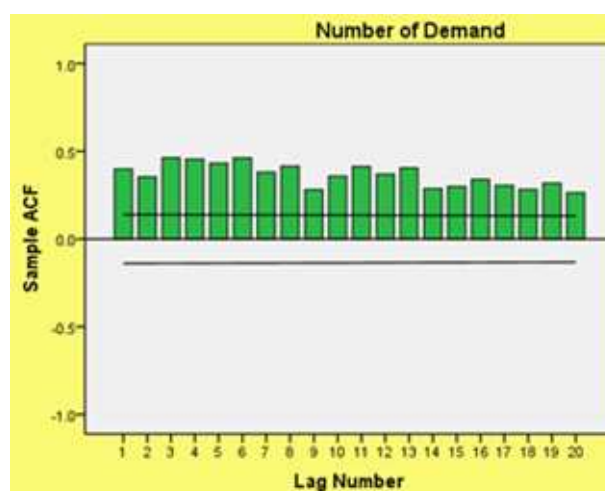
##### ۴-۱-۱. شناسایی نوع و مرتبه مدل

نمودار (۵) نمودار داده‌های اولیه سری مورد استفاده در فرایند آریما جهت پیش‌بینی تقاضای سی‌تی‌اسکن را نشان می‌دهد.



نمودار ۵. سری زمانی تقاضای سی تی اسکن

ابتدا مانایی (ایستایی) سری سنجیده می شود. نمودار (۶) تابع خودهمبستگی سری را برای ۲۰ وقفه (Lag) نشان می دهد. همان طور که در نمودار (۸) نشان می دهد مقادیر تابع خود همبستگی تا ۲۰ وقفه دارای نوعی موج سینوسی (افزایشی کاهشی) است و نشان دهنده این است که سری ممکن است ایستا نباشد. برای اطمینان بیشتر آزمون تعمیم یافته دیکی - فولر را انجام می دهیم. جدول (۱) نتیجه اجرای این آزمون را نشان می دهد.



نمودار ۶. تابع همبستگی سری

همان طور که جدول (۱) نشان می‌دهد مقدار آماره آزمون  $-2/2561$  از مقدار بحرانی  $-3/4326$  در سطح اطمینان ۹۵ درصد بیشتر است. بنابراین نمی‌توان فرض صفر را رد کرد یعنی سری دارای ریشه واحد است و ایستا نمی‌باشد. بنابراین، آزمون را یک بار دیگر اما این بار با یک مرتبه تفاضل‌گیری انجام می‌دهیم. جدول (۲) نتیجه اجرای این آزمون را نشان می‌دهد.

جدول ۱. نتایج اجرای آزمون ADF

| آماره آزمون (ADF) |           |              | (درصد) |
|-------------------|-----------|--------------|--------|
| $-2/2561$         |           |              |        |
| سطح ۱             | $-4/0050$ | مقدار بحرانی |        |
| سطح ۵             | $-3/4326$ | مقدار بحرانی |        |
| سطح ۱۰            | $-3/1401$ | مقدار بحرانی |        |

مأخذ: نتایج تحقیق.

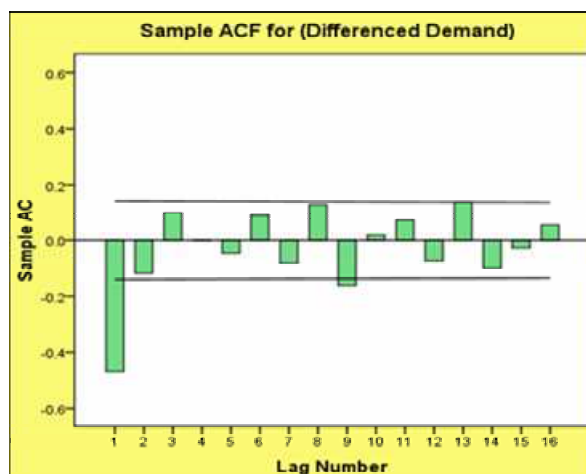
جدول ۲. نتایج اجرای آزمون ADF بعد از تفاضل‌گیری

| آماره آزمون (ADF) |           |              | (درصد) |
|-------------------|-----------|--------------|--------|
| $-2/2561$         |           |              |        |
| سطح ۱             | $-4/0050$ | مقدار بحرانی |        |
| سطح ۵             | $-3/4326$ | مقدار بحرانی |        |
| سطح ۱۰            | $-3/1401$ | مقدار بحرانی |        |

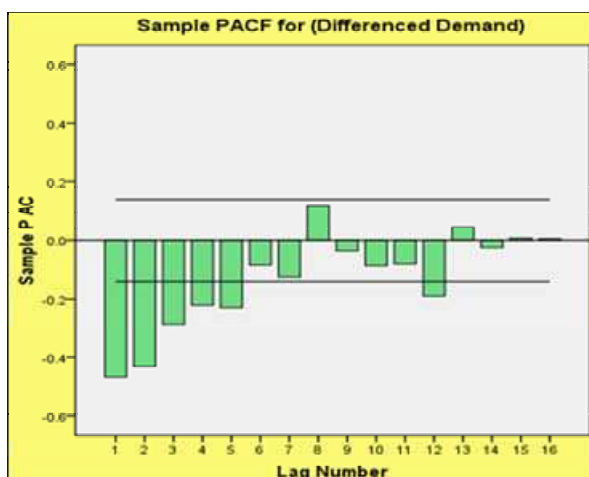
مأخذ: نتایج تحقیق.

همان طور که جدول (۲) نشان می‌دهد مقدار آماره آزمون  $-12/5437$  از مقدار بحرانی  $-3/4326$  در سطح اطمینان ۹۵ درصد کمتر است، بنابراین نمی‌توان فرض صفر را پذیرفت یعنی سری دارای ریشه واحد نیست و ایستا می‌باشد. به این ترتیب سری بعد از یک بار تفاضل‌گیری ایستا می‌شود در نتیجه مقدار پارامتر  $d=1$  می‌شود.

حال نوبت به شناسایی نوع و مرتبه مدل است. برای این کار از نمودارهای ACF و PACF سری استفاده می‌شود که اگر  $d \neq 0$  باشد باید از سری تفاضل‌گیری شده استفاده کرد. نمودار (۷) و (۸) به ترتیب نمودار تابع خودهمبستگی و تابع خودهمبستگی جزئی سری تفاضل‌گیری شده را برای ۱۶ وقفه (Lag) نشان می‌دهد.



نمودار ۷. تابع خودهمبستگی سری تفاضل‌گیری‌شده



نمودار ۸. تابع خودهمبستگی جزئی سری تفاضل‌گیری‌شده

با توجه به نمودار تابع خودهمبستگی و تابع خودهمبستگی جزئی تفاضل‌گیری‌شده برای مرتبه اتورگرسیون  $p = 1, 2, 3, 4, 5, 6$  و برای مرتبه میانگین متحرک  $q = 1, 2, 3$  جهت تعیین و تخمین پارامترهای مدل‌های اولیه در نظر گرفته شده است. به این ترتیب ۲۷ مدل برای بررسی انتخاب می‌شوند که در جدول (۳) ارائه شده است.

جدول ۳. مدل‌های مورد بررسی در فرایند ARIMA

| شماره | مدل          | شماره | مدل          | شماره | مدل          | شماره | مدل          |
|-------|--------------|-------|--------------|-------|--------------|-------|--------------|
| ۱     | ARIMA(۱,۱,۰) | ۸     | ARIMA(۰,۱,۲) | ۱۵    | ARIMA(۲,۱,۳) | ۲۲    | ARIMA(۵,۱,۱) |
| ۲     | ARIMA(۲,۱,۰) | ۹     | ARIMA(۰,۱,۳) | ۱۶    | ARIMA(۳,۱,۱) | ۲۳    | ARIMA(۵,۱,۲) |
| ۳     | ARIMA(۳,۱,۰) | ۱۰    | ARIMA(۱,۱,۱) | ۱۷    | ARIMA(۳,۱,۲) | ۲۴    | ARIMA(۵,۱,۳) |
| ۴     | ARIMA(۴,۱,۰) | ۱۱    | ARIMA(۱,۱,۲) | ۱۸    | ARIMA(۳,۱,۳) | ۲۵    | ARIMA(۶,۱,۱) |
| ۵     | ARIMA(۵,۱,۰) | ۱۲    | ARIMA(۱,۱,۳) | ۱۹    | ARIMA(۴,۱,۱) | ۲۶    | ARIMA(۶,۱,۲) |
| ۶     | ARIMA(۶,۱,۰) | ۱۳    | ARIMA(۲,۱,۱) | ۲۰    | ARIMA(۴,۱,۲) | ۲۷    | ARIMA(۶,۱,۳) |
| ۷     | ARIMA(۰,۱,۱) | ۱۴    | ARIMA(۲,۱,۲) | ۲۱    | ARIMA(۴,۱,۳) |       |              |

مأخذ: نتایج تحقیق.

## ۴-۲. تخمین و آزمون پارامترهای مدل‌های تعیین‌شده

در اینجا مدل‌ها به دو قسمت تقسیم شده‌اند. مدل‌های رد نشده (جدول ۴) و مدل‌های رد شده (جدول ۵).

جدول ۴. مدل‌های پذیرفته شده

| ضرایب    |          |          |          |          |          | مدل          |
|----------|----------|----------|----------|----------|----------|--------------|
| $\phi_1$ | $\phi_5$ | $\phi_4$ | $\phi_3$ | $\phi_2$ | $\phi_1$ |              |
|          |          |          |          |          | -۰/۴۶۸۴  | ARIMA(۱,۱,۰) |
|          |          |          |          |          | ۰/۰۰۰۰   | p-value      |
|          |          |          |          | -۰/۴۳۹۹  | -۰/۶۸۰۷  | ARIMA(۲,۱,۰) |
|          |          |          |          | ۰/۰۰۰۰   | ۰/۰۰۰۰   | p-value      |
|          |          |          | -۰/۳۰۷۴  | -۰/۶۵۶۹  | -۰/۸۱۹۶  | ARIMA(۳,۱,۰) |
|          |          |          | ۰/۰۰۰۰   | ۰/۰۰۰۰   | ۰/۰۰۰۰   | p-value      |
|          |          | ۰/۲۴۴۲   | -۰/۵۱۲۶  | -۰/۸۱۷۹  | -۰/۸۹۵۵  | ARIMA(۴,۱,۰) |
|          |          | ۰/۰۰۰۶   | ۰/۰۰۰۰   | ۰/۰۰۰۰   | ۰/۰۰۰۰   | p-value      |
|          | -۰/۲۵۹۸  | -۰/۴۸۱۹  | -۰/۷۲۴۹  | -۰/۹۴۹۵  | -۰/۹۵۹۱  | ARIMA(۵,۱,۰) |
|          | ۰/۰۰۰۳   | ۰/۰۰۰۰   | ۰/۰۰۰۰   | ۰/۰۰۰۰   | ۰/۰۰۰۰   | p-value      |
| ۰/۸۸۰۶   |          |          |          |          |          | ARIMA(۰,۱,۱) |
| ۰/۰۰۰۰   |          |          |          |          |          | p-value      |
| ۰/۸۱۰۸   |          |          |          | -۰/۲۲۲۳  | -۰/۱۷۸۶  | ARIMA(۲,۱,۱) |
| ۰/۰۰۰۰   |          |          |          | ۰/۰۰۵۳   | ۰/۰۳۰۲   | p-value      |

مأخذ: نتایج تحقیق.

این مدل‌ها دو شرط لازم برای پذیرفته شدن را دارند یعنی هم ضرایب مدل‌ها در بازه [۱ و -۱] قرار دارند و هم ضرایب آنها معنادار هستند، زیرا p-value آنها از ۰/۰۵ کوچکتر است یعنی با ۹۵ درصد اطمینان می‌توان گفت که ضرایب فوق مخالف صفر می‌باشند.

در جدول (۵) برای برخی از مدل‌ها ضرایب به دست آمده خارج از بازه [۱ و -۱] قرار دارند، پس نمی‌توان مدل را پذیرفت و برای تعداد دیگری از مدل‌ها ضرایب به دست آمده در سطح اطمینان ۹۵ درصد معنادار نیستند، زیرا p-value آنها از ۰/۰۵ بزرگتر است و این نشان‌دهنده این است که با ۹۵ درصد اطمینان می‌توان گفت که ضرایب فوق مساوی با صفر می‌باشند و چون این ضرایب مربوط به مرتبه‌های پایانی مدل در فرایند AR یا MA هستند نمی‌توان مدل را پذیرفت.

جدول ۵. مدل‌های پذیرفته نشده

| ضرایب       |             |             |             |             |             |             |             |             | مدل          |
|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|--------------|
| $\varphi_3$ | $\varphi_2$ | $\varphi_1$ | $\varphi_6$ | $\varphi_5$ | $\varphi_4$ | $\varphi_3$ | $\varphi_2$ | $\varphi_1$ |              |
|             | -۱/۱۷۰۸     | ۱/۱۷۰۸      |             |             |             |             |             |             | ARIMA(۰,۲,۱) |
|             | ۰/۰۱۰۴      | ۰/۰۰۰۰      |             |             |             |             |             |             | p-value      |
| -۰/۱۸۰۵     | ۰/۰۳۸۳      | ۰/۹۹۶۷      |             |             |             |             |             |             | ARIMA(۰,۱,۳) |
| ۰/۰۷۷۵      | ۰/۶۹۴۷      | ۰/۰۰۰۰      |             |             |             |             |             |             | p-value      |
|             |             | ۰/۸۶۲۷      |             |             |             |             | -۰/۱۰۷۸     |             | ARIMA(۱,۱,۱) |
|             |             | ۰/۰۰۰۰      |             |             |             |             | ۰/۱۷۵۴      |             | p-value      |
|             | ۰/۷۲۸۴      | ۰/۰۵۴۷      |             |             |             |             | -۰/۸۲۰۲     |             | ARIMA(۲,۱,۱) |
|             | ۰/۶۵۴۶      | ۰/۹۷۶۹      |             |             |             |             | ۰/۶۶۴۵      |             | p-value      |
| -۰/۱۸۹۹     | ۰/۱۳۸۴      | ۰/۸۹۱۸      |             |             |             |             | -۰/۱۰۷۰     |             | ARIMA(۱,۱,۳) |
| ۰/۰۲۷۳      | ۰/۷۲۷۴      | ۰/۰۲۲۹      |             |             |             |             | ۰/۷۸۷۷      |             | p-value      |
|             | -۰/۲۷۴۹     | ۱/۱۱۹۲      |             |             |             |             | -۰/۱۹۶۶     | ۰/۱۱۸۰      | ARIMA(۲,۱,۲) |
|             | ۰/۲۶۴۰      | ۰/۰۰۰۱      |             |             |             |             | ۰/۰۲۶۳      | ۰/۶۷۰۸      | p-value      |
| ۰/۷۱۱۴      | -۱/۹۲۳۳     | ۲/۱۱۷۳      |             |             |             |             | -۰/۷۲۶۳     | ۱/۰۴۲۵      | ARIMA(۲,۱,۳) |
| ۰/۰۰۰۰      | ۰/۰۰۰۰      | ۰/۰۰۰۰      |             |             |             |             | ۰/۰۰۰۰      | ۰/۰۰۰۰      | p-value      |
|             |             | ۰/۷۹۱۳      |             |             |             |             | -۰/۰۵۵۴     | -۰/۲۴۶۶     | ARIMA(۳,۱,۱) |
|             |             | ۰/۰۰۰۰      |             |             |             |             | ۰/۵۲۸۵      | ۰/۰۰۶۴      | p-value      |
|             | ۰/۸۰۳۹      | -۰/۱۸۲۴     |             |             |             |             | -۰/۲۳۵۱     | -۰/۴۰۰۲     | ARIMA(۳,۱,۲) |
|             | ۰/۰۰۰۰      | ۰/۰۰۰۷      |             |             |             |             | ۰/۰۰۲۹      | ۰/۰۰۲۴      | p-value      |
| -۰/۳۰۳۳     | ۰/۸۴۱۳      | ۰/۱۵۱۲      |             |             |             |             | -۰/۱۹۸۸     | -۰/۰۵۰۸     | ARIMA(۳,۱,۳) |

ادامه جدول ۵.

| ضرایب    |          |          |          |          |          |          |          |          | مدل          |
|----------|----------|----------|----------|----------|----------|----------|----------|----------|--------------|
| $\Phi_3$ | $\Phi_2$ | $\Phi_1$ | $\Phi_6$ | $\Phi_5$ | $\Phi_4$ | $\Phi_3$ | $\Phi_2$ | $\Phi_1$ |              |
| ۰/۱۶۰۶   | ۰/۰۰۰۰   | ۰/۵۵۵۹   |          |          |          | ۰/۰۲۳۵   | ۰/۸۵۹۳   | ۰/۰۰۱۰   | p-value      |
|          |          | ۰/۷۶۱۴   |          |          | -۰/۰۵۶۰  | -۰/۰۸۹۷  | -۰/۲۸۵۹  | -۰/۲۴۱۵  | ARIMA(۴,۱,۱) |
|          |          | ۰/۰۰۰۰   |          |          | ۰/۵۳۷۸   | ۰/۳۹۶۶   | ۰/۰۱۰۴   | ۰/۰۳۱۷   | p-value      |
|          | -۰/۷۹۴۸  | ۰/۱۹۷۰   |          |          | -۰/۰۴۲۰  | -۰/۲۸۷۲  | -۰/۴۳۰۳  | -۱/۱۵۴۸  | ARIMA(۱,۲,۴) |
|          | ۰/۰۰۰۰   | ۰/۰۰۲۸   |          |          | ۰/۶۲۸۸   | ۰/۰۵۵۵   | ۰/۰۰۷۰   | ۰/۰۰۰۰   | p-value      |
| -۰/۲۳۱۸  | ۰/۸۳۸۸   | ۰/۰۷۶۶   |          |          | ۰/۰۱۴۶   | -۰/۱۹۱۴  | -۰/۱۳۸۴  | -۰/۸۹۲۶  | ARIMA(۴,۱,۳) |
| ۰/۵۸۲۷   | ۰/۰۰۰۰   | ۰/۸۸۷۴   |          |          | ۰/۹۲۷۸   | ۰/۴۶۳۹   | ۰/۸۲۰۹   | ۰/۰۹۵۵   | p-value      |
|          |          | ۰/۶۸۷۵   |          | -۰/۰۵۹۹  | -۰/۱۱۴۵  | -۰/۱۷۱۴  | -۰/۳۶۲۵  | ۳۲۰۱/-۰  | ARIMA(۵,۱,۱) |
|          |          | ۰/۰۰۰۰   |          | ۰/۵۳۴۰   | ۰/۳۳۹۳   | ۰/۲۵۲۲   | ۰/۰۱۴۴   | ۰/۰۳۳۰   | p-value      |
|          | ۰/۶۳۲۳   | -۰/۱۲۳۷  |          | -۰/۱۱۱۵  | -۰/۱۸۳۰  | -۰/۳۸۶۶  | -۰/۵۳۸۹  | -۱/۱۰۸۹  | ARIMA(۵,۱,۲) |
|          | ۰/۰۰۰۰   | ۰/۴۸۱۱   |          | ۰/۲۲۰۹   | ۰/۲۶۱۶   | ۰/۰۴۰۹   | ۰/۰۰۵۹   | ۰/۰۰۰۰   | p-value      |
| -۰/۲۸۹۲  | ۰/۸۲۳۶   | ۰/۱۱۸۹   |          | -۰/۰۴۴۵  | -۰/۰۲۸۳  | -۰/۲۱۲۶  | ۱۱۳۱/-۰  | -۰/۸۴۶۸  | ARIMA(۵,۱,۳) |
| ۰/۲۰۶۴   | ۰/۰۰۰۰   | ۰/۶۸۸۶   |          | ۰/۶۲۶۳   | ۰/۸۶۵۶   | ۰/۲۷۳۲   | ۰/۷۴۱۷   | ۰/۰۰۴۱   | p-value      |
|          |          |          | -۰/۱۰۰۱  | -۰/۳۴۳۶  | -۰/۵۶۴۸  | -۰/۷۸۵۴  | -۰/۹۸۹۷  | -۰/۹۸۱۵  | ARIMA(۶,۱,۰) |
|          |          | ۰/۱۶۵۹   | ۰/۰۰۰۶   | ۰/۰۰۰۰   | ۰/۰۰۰۰   | ۰/۰۰۰۰   | ۰/۰۰۰۰   | ۰/۰۰۰۰   | p-value      |
|          | ۰/۸۱۰۶   | ۰/۱۰۴۸   | ۰/۰۳۴۶   | ۰/۰۰۵۵   | -۰/۰۳۸۷  | -۰/۲۳۵۴  | -۰/۱۹۴۱  |          | ARIMA(۶,۱,۱) |
|          | ۰/۰۰۰۰   | ۰/۲۸۱۳   | ۰/۷۷۶۴   | ۰/۹۷۱۳   | ۰/۸۲۵۶   | ۰/۱۶۷۸   | ۰/۲۵۵۰   |          | p-value      |
|          | ۰/۵۴۶۵   | ۰/۰۳۳۶   | ۰/۰۶۱۹   | -۰/۰۴۳۶  | -۰/۱۱۳۸  | -۰/۳۰۲۴  | -۰/۴۶۸۲  | -۰/۹۵۰۲  | ARIMA(۶,۱,۲) |
|          | ۰/۰۰۱۴   | ۰/۸۹۳۵   | ۰/۵۶۴۵   | ۰/۸۰۷۰   | ۰/۶۱۷۸   | ۰/۲۴۸۸   | ۰/۰۶۸۳   | ۰/۰۰۰۳   | p-value      |
| ۰/۷۶۵۱   | ۰/۳۱۸۰   | -۰/۹۲۲۰  | ۰/۰۲۶۰   | -۰/۱۰۸۲  | -۰/۴۰۱۴  | -۰/۷۳۶۴  | -۱/۵۴۸۵  | -۱/۸۳۹۶  | ARIMA(۶,۱,۳) |
| ۰/۰۰۰۰   | ۰/۰۴۲۴   | ۰/۰۰۰۰   | ۰/۷۸۵۴   | ۰/۶۲۲۸   | ۰/۲۲۲۷   | ۰/۰۳۹۸   | ۰/۰۰۰۰   | ۰/۰۰۰۰   | p-value      |

مأخذ: نتایج تحقیق.

پس از تعیین مدل‌های قابل قبول نوبت به محاسبه معیار ذکر شده برای مدل‌های پذیرفته شده می‌رسد که این معیارها در جدول (۶) محاسبه شده است.

جدول ۶. معیارهای انتخاب محاسبه شده

| MAPE    | SSE     | FPE    | SBC    | AIC    | مدل          |
|---------|---------|--------|--------|--------|--------------|
|         |         |        |        |        |              |
| ۳۶/۹۶۵۹ | ۸۵/۶۳۶۴ | ۰/۴۳۳۱ | ۲/۰۰۶  | ۱/۹۸۹۴ | ARIMA(۱,۱,۰) |
| ۳۴/۶۵۴۴ | ۶۸/۸۸۳۳ | ۰/۳۵۵۱ | ۱/۸۱۹۷ | ۱/۷۸۶۸ | ARIMA(۲,۱,۰) |
| ۳۳/۲۷۰۱ | ۶۲/۲۲۲۲ | ۰/۳۲۹۶ | ۱/۷۴۹۷ | ۱/۷۰۰۲ | ARIMA(۳,۱,۰) |
| ۳۲/۳۵۳۳ | ۵۸/۴۳۷۳ | ۰/۳۱۸۲ | ۱/۷۱۸۹ | ۱/۶۵۲۷ | ARIMA(۴,۱,۰) |
| ۳۰/۹۰۸۰ | ۵۴/۴۶۱۰ | ۰/۳۰۶۵ | ۱/۶۸۰۶ | ۱/۵۹۷۶ | ARIMA(۵,۱,۰) |
| ۳۲/۱۶۰۲ | ۵۷/۱۶۷۵ | ۰/۲۹۸۵ | ۱/۵۹۶۸ | ۱/۵۶۰۶ | ARIMA(۰,۱,۱) |
| ۳۰/۷۱۰۸ | ۵۴/۱۳۱۵ | ۰/۲۹۲۹ | ۱/۶۰۵۱ | ۱/۵۵۵۸ | ARIMA(۲,۱,۱) |

مأخذ: نتایج تحقیق.

همان طور که ذکر شد برای انتخاب بهترین مدل از روش TOPSIS استفاده شده است. جدول (۷) رتبه‌بندی نهایی مدل‌ها را نشان می‌دهد. به این ترتیب، با توجه به جدول (۷) مدل ARIMA(2,1,1) احتمالاً مناسب‌ترین مدل در بخش آریما است.

جدول ۷. رتبه‌بندی نهایی مدل‌ها با استفاده از TOPSIS

| رتبه | مدل          | امتیاز هر گزینه |
|------|--------------|-----------------|
| ۱    | ARIMA(۲,۱,۱) | ۰/۹۹۳۸          |
| ۲    | ARIMA(۵,۱,۰) | ۰/۹۱۳۶          |
| ۳    | ARIMA(۰,۱,۱) | ۰/۹۱۲۹          |
| ۴    | ARIMA(۴,۱,۰) | ۰/۸۰۹۹          |
| ۵    | ARIMA(۳,۱,۰) | ۰/۷۰۷۱          |
| ۶    | ARIMA(۲,۱,۰) | ۰/۵۱۰۷          |
| ۷    | ARIMA(۱,۱,۰) | ۰               |

مأخذ: نتایج تحقیق.

#### ۲-۴. بررسی تشخیصی (آزمون پسماندها)

همان طور که در بخش سوم گفته شد، این کار در اینجا توسط آزمون Ljung-Box انجام می‌شود. آزمون برای مدل ARIMA(2,1,1) برای سه لگ ۱۲، ۱۸ و ۲۴ انجام شده است (جدول ۸).

جدول ۸. نتایج آزمون Ljung-Box

| لگ | آماره Q | درجه آزادی | مقدار بحرانی | p-value |
|----|---------|------------|--------------|---------|
| ۱۲ | ۱۵/۶۶۸۷ | ۹          | ۱۶/۹۱۹۰      | ۰/۰۷۴۱  |
| ۱۸ | ۲۲/۶۴۷۹ | ۱۵         | ۲۴/۹۹۸۵      | ۰/۰۹۱۹  |
| ۲۴ | ۳۱/۸۳۵۴ | ۲۱         | ۳۲/۶۷۰۷      | ۰/۰۶۰۸  |

مأخذ: نتایج تحقیق.

همان طور که جدول (۸) نشان می‌دهد در هر ۳ لگ مورد آزمون، مقدار آماره آزمون کمتر از مقدار بحرانی است ( $p\text{-value} > 0/05$ ) پس فرض صفر رد نشده و نشان‌دهنده تصادفی بودن پسماندها می‌باشد. به این ترتیب نیازی به بررسی مدل‌های دیگر نیست و مدل فوق به عنوان بهترین مدل در بخش آرما انتخاب می‌شود. مدل نهایی به صورت زیر می‌باشد:

$$(1 - 0.1786B - 0.2223B^2)(1 - B)X_t = (1 + 0.8108B)Z_t \quad (12)$$

## ۳-۴. یافته‌های مربوط به شبکه عصبی مصنوعی

در این بخش شبکه یک بار با یک لایه پنهان و یک بار هم با دو لایه پنهان بر اساس اندازه جمعیت‌های گفته شده در بخش سوم هر یک ۵۰ بار آموزش داده شد که نتایج آن در جدول (۹) ارائه شده است.

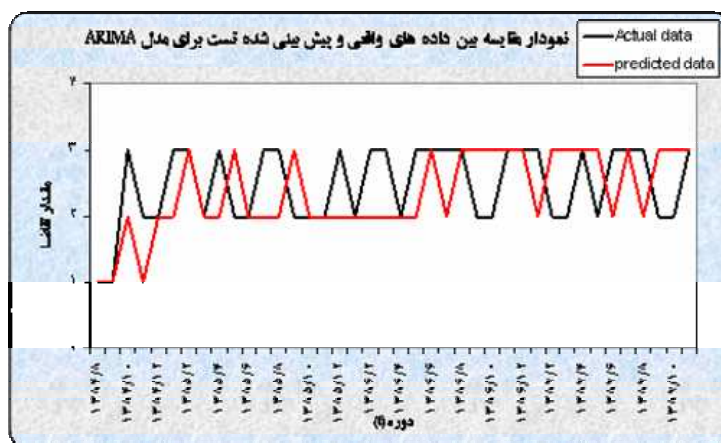
جدول ۹. نتایج اجرای شبکه عصبی مصنوعی

| اندازه جمعیت | شبکه با یک لایه پنهان |         |                          | شبکه با دو لایه پنهان |         |                          |
|--------------|-----------------------|---------|--------------------------|-----------------------|---------|--------------------------|
|              | Mse آموزش             | Mse تست | متوسط زمان آموزش (دقیقه) | Mse آموزش             | Mse تست | متوسط زمان آموزش (دقیقه) |
| ۲۰           | ۰/۲۴۱۵                | ۰/۱۶۷۶  | ۲/۲                      | ۰/۲۲۹۷                | ۰/۱۶۲۱  | ۴/۶                      |
| ۵۰           | ۰/۲۳۹۷                | ۰/۱۶۸۷  | ۶/۷                      | ۰/۲۲۹۵                | ۰/۱۶۰۲  | ۷/۶                      |
| ۸۰           | ۰/۲۳۴۷                | ۰/۱۵۹۹  | ۱۰/۶۵                    | ۰/۲۲۶۴                | ۰/۱۵۸۷  | ۱۲/۲۱                    |
| ۱۰۰          | ۰/۲۲۹۸                | ۰/۱۵۴۴  | ۷۶/۱۶                    | ۰/۲۲۳۸                | ۰/۱۵۳۱  | ۱۷/۶۵                    |
| ۱۲۰          | ۰/۲۲۱۸                | ۰/۱۵۳۹  | ۲۱/۸۸                    | ۰/۲۱۸۷                | ۰/۱۵۰۱  | ۲۱/۸۷                    |
| ۱۵۰          | ۰/۲۱۹۹                | ۰/۱۵۰۲  | ۳۵/۸۷                    | ۰/۲۰۸۶                | ۰/۱۴۶۳  | ۳۸/۲۱                    |
| ۱۸۰          | ۰/۲۲۹۰                | ۰/۱۵۶۴  | ۴۲/۵۳                    | ۰/۲۰۹۹                | ۰/۱۴۸۷  | ۴۷/۸۹                    |

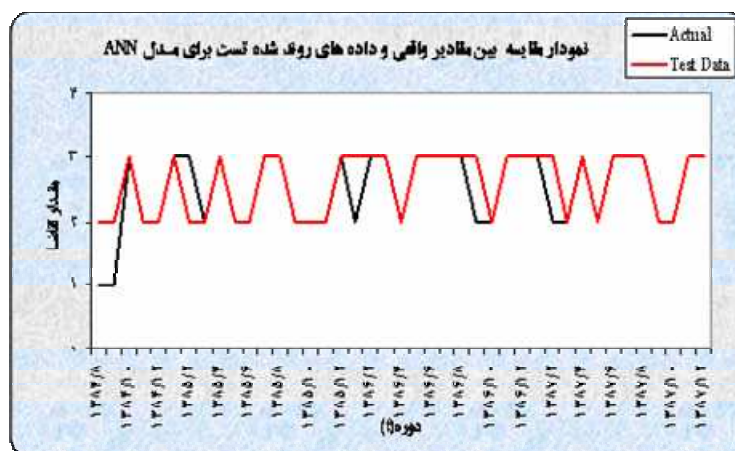
مأخذ: نتایج تحقیق.

همان طور که جدول (۹) نشان می دهد در قسمت شبکه عصبی مصنوعی، شبکه با دو لایه پنهان با اندازه جمعیت ۱۵۰ عملکرد بهتری نسبت به سایر موارد دارد و به عنوان شبکه بهینه در پیش بینی تقاضای دستگاه سی تی اسکن انتخاب می شود.

نمودار (۹) مقایسه بین اعداد واقعی و پیش بینی شده داده های تست برای مدل ARIMA و نمودار (۱۰) مقایسه مقادیر روند شده داده های تست و داده های واقعی را توسط مدل شبکه عصبی نشان می دهد.



نمودار ۹. مقایسه بین اعداد واقعی و پیش بینی شده داده های تست برای مدل آریم



نمودار ۱۰. مقایسه بین مقادیر واقعی و پیش بینی شده داده های تست برای مدل ANN

جدول (۱۰) مقادیر معیارهای سنجش خطای خروجی‌های هر دو مدل انتخابی را برای داده‌های تست برای دو معیار MSE و MAPE نشان می‌دهد.

جدول ۱۰. معیارهای سنجش خطا برای دو مدل برای داده‌های تست

| معیارها |        | مدل          |
|---------|--------|--------------|
| MAPE    | MSE    |              |
| ۹/۳۴۲   | ۰/۱۴۶۳ | ANN          |
| ۲۲/۷۶۴  | ۰/۵۶۱۰ | ARIMA(۲،۱،۱) |

مأخذ: نتایج تحقیق.

بر اساس سیستم پیشنهادشده نتایج دو مدل باید مقایسه شود و هر مدلی که خطای کمتری داشته باشد برای پیش‌بینی تقاضای دستگاه سی‌تی‌اسکن انتخاب شود. با توجه به نتایج جدول (۱۰) مدل شبکه عصبی دارای خطای کمتری است و مقادیر پیش‌بینی‌شده توسط آن بسیار نزدیکتر از مقادیر پیش‌بینی‌شده توسط مدل آریما به مقادیر واقعی است و به عنوان مدل بهینه انتخاب می‌شود.

##### ۵. نتیجه‌گیری

این مقاله یک سیستم مقایسه‌ای پیش‌بینی را برای پیش‌بینی تقاضای تجهیزات پزشکی بر اساس روش‌های آریمای تک متغیره و شبکه‌های عصبی چندلایه پیشخور توسعه داده است که برای پیش‌بینی تقاضای سایر تجهیزات پزشکی دیگر نیز قابل استفاده است. سیستم پیشنهاد شده قصد دارد تا مدل‌سازی شبکه عصبی و آریما را با کاری مستدل برای تخمین تقاضا بکار برده و مقایسه کند. بنابراین، یک رویه دو قسمتی که روش‌های شبکه عصبی و آریما را بکار می‌برد تا پیش‌بینی‌های دقیق‌تری حاصل شود. مورد دیگر اینکه پیش‌بینی در قسمت‌های مختلف بهداشت و درمان اگر به صورت سالانه انجام شود این امکان را فراهم می‌آورد تا بتوان از متغیرهای کلان ملی مرتبط با این بخش برای پیش‌بینی استفاده کرد. دستگاه سی‌تی‌اسکن جبرای سنجش تقاضای آن انتخاب شد. ارزیابی نتایج نشان می‌دهد که روش شبکه عصبی پیش‌بینی دقیق‌تری را انجام می‌دهد.

## منابع

- Atiya, A. F.** (2001), "Bankruptcy Prediction for Credit Risk Using Neural Networks: A Survey and New Results", *IEEE Transaction on Neural Networks*, Vol. 12, No. 4, PP. 929-935.
- Abraham, B. & J. Ledolter** (1983), *Statistical Methods for Forecasting*, New York: John Wiley & Sons.
- Aburto, L. & R. Weber** (2007), "Improved Supply Chain Management Bbased on Hybrid Demand Forecasts", *Applied Soft Computing*, Vol. 7, PP. 136-144.
- Alon, I., Qi, M. & R. J. Sadowski** (2001), " Forecasting Aaggregate Retail Sales: A Comparison of Artificial Neural Networks and Traditional Methods", *Journal of Retailing and Consumer Services*, Vol. 8, PP. 147 – 156.
- Al-Saba, T. & I. El-Amin** (1999), "Artificial Neural Nnetworks as Applied to Long Term Demand Forecasting", *Artificial Intelligence in Engineering*, Vol. 13, PP. 189-197.
- Beccali, M., Cellura, M., Lo Brano, V. & A. Marvuglia** (2004), "Forecasting Daily Urban Electric Lload Profiles Using Artificial Neural Networks", *Energy Conversion and Management*, Vol. 45, PP. 2879-2900.
- Chatfield, C.** (2004), *The Analysis of Time Series: An Introduction*, 6th Ed, London: Chapman & Hall, A CRC Press Company.
- Goldberg, D. E.** (1989), *Genetic Algorithms in Search, Optimization, and Machine Learning*, Addison-Wesley Publishing Company, Inc.
- Dhar, V. & R. Stein** (1997), *Seven Methods for Transforming Corporate Data into Business Intelligence*, NJ: Prentice-Hal.
- Dickey, D. A. & W. A. Fuller** (1979), "Distribution of the Estimators for Autoregressive time Series with a Unit Root", *Journal of the American Statistical Association*, Vol. 74, PP. 427-431.
- Du, T. C. T. & P. M. Wolfe** (1997), "Implementation of Fuzzy Logic Systems and Nneural Networks in Industry", *Computers in Industry*, Vol. 32, PP. 261-272.
- Schoneburg E.** (1990), "Stock Price Prediction Using Neural Networks: A Project Report", *Neurocomputing*, Vol. 2, PP. 17-27.
- Efendigil, T., Onut, S. & C. Kahraman** (2009), "A Decision Support System for Demand Forecasting with Artificial Neural Networks and Neuro-Fuzzy Models: A Comparative Analysis", *Expert Systems with Applications*, Vol. 36, PP. 6697 – 6707.
- Schwarz, G.** (1978), "Estimating the Ddimension of a Model", *Annals of Statistics*, Vol. 6, PP. 461- 464.
- Akaike, H.** (1970), "Statistical Predictor Identification", *Ann, Inst, Math*, Vol. 22, PP. 203 – 217.
- Akaike, H.** (1974), "A New Look at the Statistical Model Identification", *IEEE Transaction on Automatic Control*, AC-19, PP. 716-723.
- Hajmeer, M. & I. A. Basheer** (2000), "Artificial Neural Networks: Fundamentals, Computing, Design, and Application", *Journal of Microbiological Methods*, Vol. 43, PP. 3 – 31.
- Hobbs, B. F., Helman, U., Jitprapaikulsarn, S., Konda, S. & D. Maratukulam** (1998), "Artificial Neural Networks for Short-Term Energy Forecasting: Accuracy and Economic Value", *Neurocomputing*, Vol. 23, PP. 71-84.
- Huang, W., Lai, K. K., Nakamori, Y & S. Wang** (2004), "Forecasting Foreign Exchange Rates with Artificial Neural Networks: A Review", *International Journal of Information Technology & Decision Making*, Vol. 3, No. 1, PP.145-165.

- Holland J. H.** (1975), *Adaptation in Natural and Artificial Systems*, Ann Arbor MI: The University of Michigan Press.
- Scott, Armstrong, J.** (2001), *Principles of Forecasting: a Handbook for Researchers and Practitioners*, Norwell, Massachusetts: Kluwer Academic Publishers.
- Kuo, R. J. & C. K. Xue** (1998), "An Intelligent Sales Forecasting System Through Integration of Artificial Neural Network and Fuzzy Neural Network", *Computers in Industry*, Vol. 37, PP. 1–15.
- Kuo, R. J. Wu, P. & C. P. Wang** (2002), "An Intelligent Sales Forecasting System Through Integration of Artificial Neural Networks and Fuzzy Neural Networks with Fuzzy Weight Elimination", *Neural Networks*, Vol. 15, PP. 909–925.
- Fausett, L.** (1994), *Fundamental of Neural Networks*, Prentice-Hall.
- Ellram, L.M.** (1991), "Supply Chain Management: The Industrial Organization Perspective", *International Journal of Physical Distribution and Logistics Management*, Vol. 21, No.1, PP.13–22.
- Lachtermacher, G. & J. D. Fuller** (1995), "Backpropagation in Time Series Forecasting", *Journal of Forecasting*, Vol. 14, PP. 381–393.
- Law, R. & N. Au** (1999), "A Nneural Network Model to Forecast Japanese Demand for Travel to Hong Kong", *Tourism Management*, Vol. 20, PP. 89–97.
- Ljung, G. M. & G. E. P. Box** (1978), "On a Measure of Lack of Fit in Time Series Models", *Biometrika*, Vol. 65, PP. 297–30
- Luxhoj, J. T., Riis, J. O. & B. Stensballe** (1996), "A Hybrid Econometric-Nneural Network Modeling Approach for Sales Forecasting", *International Journal of Production Economics*, Vol. 43, PP. 175–192.
- Palmer, A., Montano, J. J. & A. Sese** (2006), "Designing an Artificial Neural Network for Forecasting Tourism Time Series", *Tourism Management*, Vol. 27, PP. 781–790.
- Palmer, A., Montano, J. J. & A. Sese** (2006), "Designing an Artificial Neural Network for Forecasting Tourism Time Series", *Tourism Management*, Vol. 27, PP. 781–790.
- Porter, M.E. & S. Stern** (2001), "Innovation: Location Matters", *MIT Sloan Management Review*, Vol. 42, No. 4, PP. 28–36.
- Kuo, R. J.** (2001), "A Sales Forecasting System Based on Fuzzy Nneural Network with Initial Weights Generated by Genetic Algorithm", *European Journal of Operational Research*, Vol. 129, PP. 496–517.
- Law, R.** (2000), "Back-Propagation Learning in Improving the Accuracy of Neural Network-Based Tourism Demand Forecasting", *Tourism Management*, Vol. 21, PP. 331–340.
- Haykin, S.** (1994), *Neural Networks: A Comprehensive Foundation*, New York: Macmillan.
- Sivanandam, S.N. & S.N. Deepa** (2008), *Introduction to Genetic Algorithms*, Berlin Heidelberg: Springer-Verlag.
- Sozen, A., Arcaklioglu, E. & M. Ozkaymak** (2005), "Turkey's Net Energy Consumption", *Applied Energy*, Vol. 81, No. 2, PP. 209–221.
- Srinivasan, D., Liew, A.C. & C. S. Chang** (1994), "A Neural Nnetwork Short-Term Load Forecaster", *Electric Power Systems Research*, Vol. 28, PP. 227–234.
- Sun, Z. L., Choi, T. M., Au, K. F. & Y. Yu** (2008), "Sales Forecasting Using Extreme Learning Machine with Applications in Fashion Retailing", *Decision Support Systems*, Vol. 46, PP. 411 – 419.
- Tang, Z. & P. A. Fishwick** (1993), "Feedforward Neural Nets as Models for Time Series Forecasting", *ORSA Journal on Computing*, Vol. 5, No. 4, PP. 374–385.

- Tang, Z., Almeida, C. & P. A. Fishwick** (1991), "Times Series Forecasting Using Neural Networks vs. Box-Jenkins Methodology", *Simulations Councils*, PP. 303-310.
- Weigend, A.S., Huberman, B.A. & D.E. Rumelhart** (1992), *Predicting Sunspots and Exchange Rates with Connectionist Networks*, In: Casdagli, M., Eubank, S. (Eds.), *Nonlinear Modeling and Forecasting*. Addison-Wesley, Redwood City, CA, PP. 395-432.
- Zhang, G. P., Patuwo, E. P. & M. Y. Hu** (1998). "Forecasting with Artificial Neural Networks: The State of the Art", *International Journal of Forecasting*, Vol. 14, PP. 35-62.