

## بهینه‌سازی پیش‌بینی تقاضای وجه نقد دستگاه‌های خودپرداز شبکه بانکی کشور با استفاده از شبکه عصبی بازگشتی عمیق LSTM

سودابه پورذاکر عربانی<sup>۱\*</sup>، حسین ابراهیم پور کومله<sup>۲</sup>

۱- دانشجوی دکتری، دانشگاه کاشان، هوش مصنوعی، کاشان، اصفهان، ایران

۲- استادیار، دانشگاه کاشان، گروه کامپیوتر، کاشان، اصفهان، ایران

رسید مقاله: ۱۷ مهر ۱۳۹۷

پذیرش مقاله: ۱۸ خرداد ۱۳۹۸

### چکیده

یکی از مشکلات سیستم‌های بانکی، پیش‌بینی تقاضای وجه نقد خودپردازها است. پیش‌بینی صحیح می‌تواند به دلایل زیر باعث سودآوری سیستم بانکی و رضایت‌مندی مشتریان این سیستم بانکی گردد. دقت در پیش‌بینی، هدف اصلی این پژوهش است. اگر خودپردازها با کمبود وجه نقد مواجه شوند محبوبیت بانک ارایه‌دهنده این سرویس کاهش خواهد یافت و بانک با کاهش استفاده مشتریان از این سیستم مواجه خواهد شد. از طرفی دیگر اگر بانک دچار محبوس شدن وجه نقد در خودپرداز شود، با توجه به تورم در ایران، این وضعیت روی سودآوری بانک تأثیر منفی خواهد گذاشت؛ بنابراین هدف از این پژوهش، پیش‌بینی دقیق برای رفع هزینه‌های دوگانه است؛ چون اطلاعات میزان وجه نقد به صورت روزانه است، بنابراین هر خودپرداز، رفتاری به صورت سری زمانی خواهد داشت و از طرفی چون هدف ما از این پژوهش، پیش‌بینی میزان تقاضای وجه نقد همه خودپردازهاست، در نتیجه ما با داده‌هایی از نوع پل مواجه هستیم. روش‌هایی که در این تحقیق برای پیش‌بینی مورد استفاده قرار گرفته است، عبارتند از: روش آماری، روش شبکه عصبی MLP و شبکه عصبی بازگشتی عمیق LSTM. نتایج حاصل از این سه روش را سپس مورد مقایسه قرار می‌دهیم و نشان می‌دهیم روش شبکه عصبی بازگشتی عمیق LSTM دارای بالاترین دقت است.

**کلمات کلیدی:** پیش‌بینی آماری، پیش‌بینی هوشمند، شبکه عصبی مصنوعی MLP، شبکه عصبی بازگشتی عمیق LSTM، تقاضاهای وجه نقد دستگاه‌های خودپرداز.

### ۱ مقدمه

بر اساس آمار موجود در بانک مرکزی جمهوری اسلامی ایران بیش از ۴۵۰ میلیون تراکنش سالانه برای دستگاه‌های خودپرداز در کشور ثبت می‌گردد. با در نظر گرفتن این مقدار تراکنش برای یک دوره یک ساله، ما با

\* عهده دار مکاتبات

آدرس الکترونیکی: Soodabeharabani@gmail.com

داده‌های بزرگ مواجه هستیم. این حجم داده از دوجنبه حایز اهمیت است. جنبه اول از دیدگاه اقتصادی، هر بانک به عنوان یک بنگاه اقتصادی، به دنبال حداکثر کردن سود خود و همچنین به دنبال حداقل نمودن هزینه‌های تراکنش‌های مالی خودش است. بنابراین به طور مستقیم، تقاضای وجه نقد بر روی سودآوری هر بانک تأثیرگذار است و به طور غیرمستقیم، وجود و عدم وجود وجه نقد در هر خودپرداز بر روی شهرت تجاری بانک تأثیر گذاشته و در پی آن بر روی سودآوری بانک تأثیر به‌سزایی خواهد داشت. از سوی دیگر، با در نظر گرفتن شکل‌گیری داده‌های سری زمانی با حجم زیاد، اهمیت ایجاد یک مدل شبکه عصبی بازگشتی عمیق LSTM<sup>1</sup>، برای بهینه‌سازی پیش‌بینی وجود دارد به طوری که توانایی تشخیص وابستگی‌های زمانی بیش از چند صد مرحله و همچنین توانایی لازم برای تشخیص تعاملات در هم پیچیده داده‌های مالی و اقتصادی را داشته باشد [۱].

یکی از موضوعاتی که امروزه در نظام بانکی اهمیت دارد، دقت در پیش‌بینی میزان وجه نقد قابل برداشت از دستگاه‌های خودپرداز می‌باشد. میزان تقاضای وجه نقد به اطلاعات گذشته میزان برداشت و تکمیل موجودی دستگاه‌های خودپرداز بستگی دارد. اگر این پیش‌بینی درست انجام نشود و یا اینکه اصلاً روش درست و علمی برای این پیش‌بینی وجود نداشته باشد از دو جنبه بانک مورد نظر با ضررهایی مواجه می‌شود؛ یعنی اگر میزان پول بیش از نیاز باشد بخش قابل توجهی از وجه نقد بدون استفاده باقی می‌ماند و با توجه به تورم قابل ملاحظه‌ای که در کشور وجود دارد بی‌استفاده ماندن پول برای آن سیستم بانکی هزینه‌هایی را برای کم ارزش‌تر شدن پول خواهد داشت و از سوی دیگر کم بودن پول نقد در دستگاه خودپرداز باعث نارضایتی مشتریان این سیستم بانکی و کاهش شهرت تجاری بانک مربوط خواهد شد که این حالت نیز برای سیستم بانکی به دلیل از دست دادن [۲]. بنابراین ما به دنبال روشی هستیم تا میزان مطلوب وجه نقد بخشی از مشتریان خود هزینه‌هایی را خواهد داشت در دستگاه‌های خودپرداز را پیش‌بینی کنیم و باعث کاهش قابل توجهی در هزینه‌های دوگانه آن شویم. با بررسی‌هایی که انجام دادیم متوجه شدیم که در ایران هیچ روش علمی و مناسبی برای پیش‌بینی میزان تقاضای وجه نقد برای دستگاه‌های خودپرداز وجود ندارد. سیاستی که معمولاً بانک‌ها مورد استفاده قرار می‌دهند به این صورت است که بانک از طریق انعقاد قرارداد با شرکت‌های خدماتی ارائه‌دهنده این نوع سرویس، حق الزحمه ثابت و قابل توجهی را برای تکمیل موجودی دستگاه‌های خودپرداز و هزینه‌های حمل و نقل مربوط پرداخت می‌نماید [۳].

یافته‌های این تحقیق که بهینه‌سازی پیش‌بینی تقاضای وجه نقد دستگاه‌های خودپرداز به کمک روش هوشمند شبکه عصبی مصنوعی بازگشتی عمیق LSTM است، می‌تواند برای مدیران نظام بانکی در جهت کاهش هزینه‌های مستقیم و غیرمستقیم کمبود یا مازاد وجه نقد در دستگاه‌های خودپرداز مفید واقع شود. هزینه‌های مستقیم و غیر مستقیم تکمیل موجودی دستگاه‌های خودپرداز عبارتند از: هزینه درج، هزینه تنظیم مجدد، هزینه حذف کردن.

<sup>1</sup> Long Short Term Memory

منظور از هزینه درج، هزینه تکمیل موجودی است که شامل هزینه‌های نیروی انسانی، هزینه حمل پول و هزینه تعمیرات دوره‌ای دستگاه خودپرداز می‌باشد. هزینه تنظیم مجدد، شامل هزینه‌های تعمیر و نگهداری بدون برنامه و هزینه فرصت از دست رفته است. منظور از هزینه فرصت از دست رفته، هزینه‌ای است که در صورت فقدان پول نقد به بانک تحمیل می‌شود که این هزینه نیز شامل دو بخش ملموس و غیرملموس است. بخش ملموس، شامل عدم استفاده از منابع درآمدی تراکنش مورد نظر و بخش غیرملموس، کاهش مراجعه بعدی به دستگاه خودپرداز و کاهش شهرت تجاری بانک ارایه‌دهنده این خدمت است. هزینه حذف شامل موارد و زمانی است که پول در دستگاه خودپرداز وجود دارد؛ ولی مراجعه برای برداشت پول صورت نمی‌گیرد. در واقع این هزینه به دلیل محبوس بودن منابع نقدی در دستگاه اتفاق می‌افتد که این هزینه متناسب با نرخ متوسط سود تضمین شده بانکی محاسبه می‌شود.

با توجه به این که داده‌های مربوط به وجه نقد در دستگاه‌های خودپرداز به صورت روزانه وجود دارد به نظر می‌رسد که رفتاری تابع زمان می‌توان از آن انتظار داشت. از طرفی دیگر چون هدف ما پیش‌بینی تقاضای وجه نقد همه دستگاه‌های خودپرداز شبکه بانکی کشور به طور همزمان است؛ بنابراین نوع داده‌های مذکور را می‌توان داده‌های پنل قلمداد کرد. حال این سوال وجود دارد که چه روشی برای پیش‌بینی دقیق وجه نقد دستگاه‌های خودپرداز می‌تواند مورد استفاده قرار گیرد، با مرور متون و ادبیات تحقیق برای پیش‌بینی داده‌های پنل دو دسته روش قابل تفکیک است:

- پیش‌بینی داده‌های پنل به صورت سنتی آماری

- پیش‌بینی داده‌های پنل به صورت هوشمند

در ابتدا لازم است که به شرح و توصیف داده‌های پنل پردازیم. داده‌های پنل، مجموعه‌ای از داده‌ها هستند که شامل چند مقطع و یک دوره زمانی است. در این تحقیق منظور از مقطع، بیانگر دستگاه‌های خودپرداز است. در حالت کلی تعداد مقطع را با  $n$  نشان می‌دهیم. دوره زمانی نیز می‌تواند روز، هفته، ماه، فصل و سال باشد که دوره زمانی استفاده شده در این جا همان روزهای سال است. طول دوره زمانی را با  $T$  نشان می‌دهیم؛ بنابراین مشاهدات را با  $X_{it}$  بیان می‌کنیم که مقطع‌ها شامل  $i = 1, 2, \dots, n$  و زمان شامل  $t = 1, \dots, T$  است؛ یعنی در کل  $nT$  مشاهده داریم [۴].

داده‌های پنل به دلیل این که هم تغییرات زمانی و هم تغییرات درون مقطعی را منعکس می‌کند می‌تواند اطلاعات بیش‌تری را شامل شود. بسیاری از نکاتی که در تحلیل سری‌های زمانی نادیده گرفته می‌شود و یا غیرقابل مشاهده هستند در تحلیل داده‌های پنل روشن می‌شوند به ویژه ناهمگنی‌هایی که غالباً در تحلیل سری زمانی از آن‌ها چشم پوشی می‌شود و غیر قابل مشاهده است، در تحلیل داده‌های پنل امکان بررسی آن‌ها فراهم می‌شود؛ یعنی در برآورد میزان برداشت آتی دستگاه‌های خودپرداز، اگر از داده‌های سری زمانی استفاده کنیم به طور ضمنی فرض کرده‌ایم که رفتار دستگاه‌های خودپرداز همگن هستند؛ اما از طرف دیگر وقتی که از داده‌های مقطعی استفاده می‌کنیم فقط به تفاوت رفتار برداشت وجه نقد دستگاه‌های خودپرداز توجه کرده‌ایم و تغییرات

زمانی آن‌ها را نادیده گرفته‌ایم. به این ترتیب داده‌های پنل امکان بررسی هر دو تغییر (مقطعی و زمانی) را فراهم می‌کند.

هدف ما در این تحقیق، ایجاد یک معماری بهینه برای شبکه عصبی بازگشتی عمیق LSTM، برای پیش‌بینی داده‌های پنل با حجم بزرگ مربوط به دستگاه‌های خود پرداز شبکه بانکی است که دارای دقت و صحت مناسبی برای این پیش‌بینی باشد.

## ۲ پیشینه تحقیق

تحقیقاتی در حوزه پیش‌بینی میزان تقاضای وجه نقد برای دستگاه‌های خودپرداز وجود دارد که در ادامه به بررسی آن‌ها خواهیم پرداخت. تدی و ان جی در پژوهش خود از شبکه عصبی<sup>۱</sup> PSECMAC استفاده کردند [۵] تا تقاضای وجه نقد روی مجموعه داده‌ای NN5 را پیش‌بینی نمایند. آن‌ها بیان کردند که این روش مقدار میانگین SMAPE<sup>۲</sup> را ۲۸/۲۷٪ از داده‌های NN5 به روش hold-out دارد.

آندراویس و همکارانش در تحقیق خود، شبکه عصبی، مدل‌های خطی و رگرسیون را با هم ترکیب کردند [۶] و با استفاده از مجموعه داده‌ای NN5، آن‌ها مقادیر مورد تقاضا را برای ۵۶ روز بعدی پیش‌بینی کردند. آن‌ها گزارش دادند که روش‌شان مقدار ۱۸/۱۹۵٪، SMAPE را تولید می‌کند.

ونکاتش و همکارانش در مقاله خود، رفتار برداشت وجه نقد دستگاه‌های خودپرداز را به الگوهای مشابه با استفاده از الگوریتم تیلور-بوتینا خوشه‌بندی کردند [۷]. برای هر خوشه آن‌ها شبکه عصبی رگرسیون عمومی (GRNN) و شبکه عصبی چند لایه Feed Forward (MLFF) و متد گروه‌بندی (GMDH) و شبکه عصبی موجی (WNN) را روی مجموعه داده‌ای NN5 به کار بردند. آن‌ها بیان کردند که روش GRNN بهترین نتیجه را به دست آورده که مقدار SMAPE آن ۱۸/۴۴٪ است. همچنین آن‌ها بیان کردند که عملکردشان بهتر از روش آندراویس و همکارانش شده است.

سیموتیس و همکارانش تکنیک‌های شبکه عصبی مصنوعی و رگرسیون پشتیبان بردار (SVR) را به کار بردند [۸] و بیان کردند که شبکه عصبی مصنوعی به طور واضحی عملکرد بهتری از روش بر پایه SVR دارد. داده‌های آن‌ها شامل ۱۵ تا دستگاه خودپرداز و مجموعه آموزشی آن‌ها شامل مقادیر تقاضای وجه نقد ۲ سال بوده است و MAPE، مقداری بین ۱۵٪ و ۱۸٪ برای شبکه عصبی مصنوعی و ۱۷٪ و ۴۰٪ برای روش رگرسیون پشتیبان بردار داشته است.

بردا و همکارانش در تحقیق خود متدی را برای پیش‌بینی تقاضای وجه نقد اگریکل بانک ارایه دادند [۹]. روش آن‌ها بر پایه سری زمانی و رگرسیون برای پیش‌بینی کردن مقدار پولی بود که روزانه باید در دستگاه خودپرداز جایگذاری شود. روش‌هایی که آن‌ها برای مدل‌بندی کردن استفاده کرده بودند روش ARIMA و

<sup>1</sup> Pseudo Self-Evolving Cerebellar Model Articulation Controller

<sup>2</sup> Symmetric Mean Absolute Percentage Error

هموار کردن توانی بود. این مدل در محیط محاسباتی آماری R پیاده سازی شد. هدف آن‌ها از این تحقیق مینیمم کردن هزینه‌های اگریکل بانک بود.

پاشکار دنده کار و همکارانش در تحقیق خود نشان دادند [۱۰] که دستگاه خودپرداز یکی از محبوب‌ترین سرویس‌های بانکی است. آن‌ها در مدل خود از ژنتیک الگوریتم برای تشخیص استراتژی مناسب برای پر کردن مجدد پول برای هر دستگاه خودپرداز استفاده کردند. آن‌ها همه تراکنش‌های انجام شده در طول سال‌های ۲۰۱۱ و ۲۰۱۲ از بانک آینده در ایران را در تحقیق خود استفاده کردند. آن‌ها دستگاه‌های خودپرداز را به سه کلاس پایین، متوسط و بالا دسته‌بندی کردند.

گاریسا و گومز یک معماری شبکه عصبی جدیدی ارائه کردند [۱۱] که از اطلاعات به دست آورده شده از شبکه عصبی بازگشتی و تجزیه موجی استفاده کردند. آن‌ها مقدار SMAPE را ۲۷٪ برای داده‌های NN5 به دست آوردند و گزارش دادند که این روش عملکرد بهتری را نسبت به شبکه عصبی Feed Forward و شبکه عصبی بازگشتی فراهم می‌کند.

اگرچه تحقیقات اخیر روی تکنیک‌های هوش محاسباتی جدید برای پیش‌بینی تقاضای وجه نقد در دستگاه‌های خودپرداز تمرکز دارد، کاکاتی کاتال و همکارانش در تحقیق خود افزایش ویژگی‌های مجموعه داده‌ای برای بهبود نتایج پیش‌بینی از مدل‌های پیش‌بینی آماری را مورد توجه قرار دادند [۱۲]. در این مقاله ۱۹ روز خاص در UK در نظر گرفته شده بود و مجموعه داده‌ای NN5 که شامل ۷۳۵ روز از ۱۱۱ تا دستگاه خودپرداز در UK بود که با روزهای تقویم به روزرسانی شده و بعد از گام پیش‌پردازش و کاربرد روش هموار کردن توانی، ۲۱/۵۷٪ میانگین SMAPE را برای ۵۶ روز پیش‌بینی به دست آوردند. آن‌ها در این مطالعه نشان دادند که نتایج خوب پیش‌بینی از طریق بهبود داده‌ها حتی اگر ما تکنیک‌های هوش محاسباتی پیچیده را به کار نبریم به دست آورده می‌شود.

پاکاج جدول و همکارانش در تحقیق خود بیان کردند [۱۳] که پیش‌بینی مخزن وجه نقد دستگاه‌های خودپرداز یک عمل پیچیده است. آن‌ها در مقاله خود پیش‌بینی تقاضای وجه نقد را روی سری‌های زمانی NN5 با استفاده از شبکه عصبی مصنوعی انجام دادند. آن‌ها مجموعه داده‌ای را از ۱۱۱ تا سری زمانی به ۱۱ تا کاهش دادند. هدف اصلی آن‌ها از این تحقیق، پیش‌بینی تقاضای وجه نقد با استفاده از شبکه عصبی بود. علاوه بر آن، همان پردازش را روی دستگاه‌های خودپرداز خوشه‌بندی شده نیز به کار بردند. میزان خطای RMS را برای اطلاعات خوشه‌بندی شده به کار بردند و میانگین آن‌ها را به دست آوردند. میزان خطای RMS نشان داد که خوشه‌بندی اطلاعات قبل از به کار بردن شبکه عصبی مصنوعی، دقت پیش‌بینی مخزن وجه نقد دستگاه‌های خودپرداز را افزایش می‌دهد.

عباسی و همکارانش به پیش‌بینی نقدینگی مورد نیاز هفت دستگاه خودپرداز بانک مهر اقتصاد خراسان رضوی پرداختند [۱۴]. آن‌ها در تحقیق خود از روش غیرخطی شبکه عصبی مصنوعی با ساختار پرسپترون چند لایه و الگوریتم پس‌انتشار خطا و روش آماری ARIMA استفاده کردند. آن‌ها نشان دادند که شبکه عصبی مصنوعی نسبت به روش ARIMA نتایج بهتری دارد.

تقوی فرد و همکارانش به تحقیق درباره افزایش رضایت مشتریان دستگاه‌های خودپرداز بانک شهر و کاهش هزینه‌های آن با به‌کارگیری سیستم کنترل موجودی دستگاه‌های خودپرداز پرداختند [۱۵]. آن‌ها افزایش میزان رضایت مشتریان از دستگاه‌های خودپرداز بانک شهر را با استفاده از پرسشنامه و توزیع آن بین ۲۰۰ نفر که از دستگاه خودپرداز استفاده کردند، قبل و بعد از پیاده‌سازی سیستم کنترل موجودی وجه نقد دستگاه‌های خودپرداز بانک شهر بررسی کردند. این پژوهش با ارایه مدل حداقل کردن مجموع هزینه‌های خواب پول و فرصت از دست رفته برای دستگاه‌های خودپرداز، به بررسی میزان رضایت مشتریان بر اساس میزان فعال‌بودن خودپردازها پرداخته است. نتایج پژوهش آن‌ها نشان داد که مدل آن‌ها قادر است نقطه سفارش مجدد، میزان سفارش پول و ترکیب انواع اسکناس را تا سقف موردنظر برای هر یک از دستگاه‌های خودپرداز ارایه کند. از مهم‌ترین یافته‌های این پژوهش می‌توان به شناسایی عوامل موثر بر سطح رضایت‌مندی مشتریان دستگاه‌های خودپرداز و همچنین مدلی برای بهینه‌سازی هزینه‌های از دست دادن مشتری و مازاد پول در خودپردازها بیان نمود. به طوری که نتایج تحقیق آن‌ها نشان داد که سطح رضایت‌مندی مشتریان ۲۵ درصد رشد داشته است و به ۸۰ درصد رسیده و هزینه‌های نگهداری خودپردازها نیز کاهش قابل ملاحظه‌ای داشته است.

حاجی مولانا و همکارانش به طراحی یک مدل پویا برای پیش‌بینی میزان تقاضای وجه نقد خودپردازهای بانک شهر تهران پرداختند [۱۶]. آن‌ها به بررسی ۲۷۲ دستگاه خودپرداز در شهر تهران برای ارزیابی رفتار دستگاه‌های خودپرداز پرداختند تا هزینه‌های نگهداری برای دستگاه‌های خودپرداز حداقل شود. نتایج پژوهش آن‌ها نشان داد که این مدل قادر است میزان پول مورد نیاز را پیش‌بینی کند و با این کار هزینه بانک شهر در این حوزه به یک دهم کاهش یابد.

### ۳ روش انجام تحقیق

در این پژوهش از سه روش برای پیش‌بینی داده‌های پنل میزان تقاضای وجه نقد دستگاه‌های خودپرداز کشور استفاده شده است که عبارتند از:

- روش آماری
  - روش شبکه عصبی مصنوعی MLP<sup>۱</sup>
  - روش شبکه عصبی مصنوعی بازگشتی عمیق LSTM
- در این تحقیق به بررسی مقایسه‌ای سه روش آماری و شبکه عصبی مصنوعی MLP و شبکه عصبی مصنوعی بازگشتی عمیق LSTM برای آنالیز و پیش‌بینی کلان داده‌های از نوع پنل و انتخاب روش بهینه برای پیش‌بینی تقاضای وجه نقد دستگاه‌های خودپرداز شبکه بانکی کشور خواهیم پرداخت.

<sup>1</sup> Multilayer Perceptron

## ۳-۱ روش آماری

با توجه به انواع مختلف داده‌ها، روش‌های سنتی آماری متنوعی برای پیش‌بینی و تخمین مدل رگرسیونی وجود دارد. انواع داده‌ها به سه دسته مقطعی، سری زمانی و ترکیبی تقسیم می‌شوند.

اگر داده‌های مساله مورد نظر؛ یعنی میزان برداشت وجه نقد دستگاه‌های خودپرداز را به صورت مقطعی در نظر بگیریم، در این صورت فقط به تفاوت‌های رفتار برداشت وجه نقد دستگاه‌های مختلف خودپرداز در یک روز توجه کرده‌ایم و تغییرات زمانی میزان برداشت وجه نقد در طول سال را نادیده گرفته‌ایم؛ بنابراین این نوع مدل داده‌ای؛ یعنی مقطعی، مدل مناسبی برای پیش‌بینی تقاضای وجه نقد دستگاه‌های خودپرداز شبکه بانکی نیست. حال اگر مدل داده‌ای سری زمانی را برای پیش‌بینی میزان تقاضای وجه نقد دستگاه‌های خودپرداز در نظر بگیریم، به این ترتیب فرض کرده‌ایم که رفتار دستگاه‌های خودپرداز همگن هستند و بر روی هم هیچ تاثیری ندارند و فقط به تغییرات زمانی میزان برداشت وجه نقد هر دستگاه خودپرداز به تنهایی توجه کرده‌ایم؛ بنابراین بهترین مدل داده‌ای برای مساله مورد نظر، مدل داده‌ای ترکیبی است؛ البته داده‌های ترکیبی نیز به دو دسته قابل تفکیک هستند، یا دارای حالت تجمیعی<sup>۱</sup> هستند و یا به صورت پنل. برای تشخیص نوع داده ترکیبی، از آزمون آماری F استفاده کردیم. این آزمون دارای دو فرضیه است:

$$H_0: a_1 = a_2 = \dots = a_n \quad \text{عرض از مبدا با هم برابر است}$$

$$H_1: a_1 \neq a_2 \neq \dots \neq a_n \quad \text{عرض از مبدا متفاوت است}$$

اگر احتمال ارزیابی شده توسط این آزمون بزرگ‌تر از ۵٪ باشد فرض صفر رد نمی‌شود و از مدل رگرسیونی داده‌های تجمیعی استفاده می‌کنیم در غیر این صورت، مدل داده‌ای به صورت پنل خواهد بود. اگر فرض صفر رد شود، در فرض یک این سوال مطرح می‌شود که مدل دارای اثرات ثابت است یا تصادفی. برای پاسخ به این سوال از آزمون آماری هاسمن استفاده می‌شود. این آزمون نیز دارای دو فرضیه است:

$$H_0: b_s = \hat{B}_s$$

$$H_1: b_s \neq \hat{B}_s$$

اگر احتمال ارزیابی شده توسط آزمون هاسمن بیش‌تر از ۵٪ باشد فرض صفر رد نمی‌شود و مدل داده‌ای دارای اثرات تصادفی است و باید از مدل رگرسیونی با اثرات تصادفی برای پیش‌بینی استفاده نماییم. اگر احتمال ارزیابی شده کوچک‌تر از ۵٪ باشد فرض صفر رد می‌شود و باید از مدل رگرسیونی با اثرات ثابت برای پیش‌بینی استفاده نماییم. حال به توصیف هر یک از مدل‌های رگرسیونی می‌پردازیم.

**مدل تجمیعی<sup>۲</sup>** مدل تجمیعی بیان‌گر آن است که اثرات فردی وجود ندارد و همه گروه‌ها یکسان هستند؛ لذا معادله رگرسیون آن به صورت زیر نوشته می‌شود:

$$Y_{it} = a + BX_{it} + \varepsilon_{it} \quad (1)$$

<sup>1</sup> Pooling

<sup>2</sup> Pool

این مدل را می‌توان با روش OLS<sup>۱</sup> برآورد نمود. به این منظور لازم است که مجموع مجذور خطاها؛ یعنی  $RSS_{pooled}$  را حداقل نماییم.

$$RSS_{pooled} = \sum_{i=1}^n \sum_{t=1}^T (Y_{it} - \hat{Y}_{it})^2 = \sum_{i=1}^n \sum_{t=1}^T (Y_{it} - \hat{a} - \hat{B} X_{it})^2 \quad (2)$$

با مشتق‌گیری رابطه بالا نسبت به  $\hat{a}$  و  $\hat{B}$  خواهیم داشت:

$$\frac{\partial RSS_{pooled}}{\partial \hat{a}} = 0 \Rightarrow \sum_{i=1}^n \sum_{t=1}^T (Y_{it} - \hat{a} - \hat{B} X_{it}) = 0 \quad (3)$$

$$\frac{\partial RSS_{pooled}}{\partial \hat{B}} = 0 \Rightarrow \sum_{i=1}^n \sum_{t=1}^T (Y_{it} - \hat{a} - \hat{B} X_{it}) X_{it} = 0 \quad (4)$$

با حل این دو معادله، تخمین‌زنده‌های تجمیعی به دست می‌آیند.

$$\hat{a}_{pooled} = \bar{Y} - \hat{B} \bar{X} \quad (5)$$

$$\hat{B}_{pooled} = \frac{\sum_i \sum_t x_{it} y_{it}}{\sum_i \sum_t x_{it}^2} \quad (6)$$

$\bar{Y}$  و  $\bar{X}$  بیانگر میانگین‌های کل می‌باشند و حروف کوچک از میانگین کل را نشان می‌دهند [۴].

$$\bar{Y} = \frac{\sum_i \sum_t Y_{it}}{nT} \quad \bar{X} = \frac{\sum_i \sum_t X_{it}}{nT} \quad (7)$$

$$x_{it} = X_{it} - \bar{X} \quad y_{it} = Y_{it} - \bar{Y} \quad (8)$$

**مدل اثرات ثابت<sup>۲</sup>** در مدل اثرات ثابت فرض می‌شود که تفاوت‌های فردی یا گروهی را می‌توان در جمله ثابت منعکس نمود. هر  $a_i$  یک ضریب مجهول است که باید برآورد گردد.  $a_i$  بیانگر اثر تمامی عواملی است که به صورت مقطعی بر  $Y_{it}$  تاثیر می‌گذارند؛ اما اثر این عوامل در طول زمان ثابت است. فرض کنید که  $Y_i$  و  $X_i$  شامل T مشاهده برای گروه iام است. در این حالت معادله (۹) را داریم:

$$Y_{it} = BX_{it} + a_i + \varepsilon_{it} \quad (9)$$

$a_i$  برای هر گروه متفاوت است. برای برآورد  $a_i$  برای هر گروه یک متغیر مجازی تعریف می‌شود؛ لذا با استفاده از متغیرهای مجازی می‌توان مدل را به صورت معادله (۱۰) نوشت:

$$Y_{it} = BX_{it} + a_1 D_1 + a_2 D_2 + \dots + a_n D_n + \varepsilon_{it} \quad (10)$$

به عنوان مثال  $D_1$  برای گروه ۱ برابر ۱ و برای گروه‌های دیگر برابر صفر است. برای گروه دوم نیز  $D_2$  برابر ۱ است و برای سایر گروه‌ها برابر با صفر است. حال ضرایب مدل را می‌توان با حداقل نمودن مجموع مجذور خطاها به دست آورد. چون این مدل معروف به مدل اثرات ثابت است؛ لذا RSS آن را با  $RSS_{FE}$  نشان می‌دهیم.

<sup>۱</sup> Ordinary least square

<sup>۲</sup> Fixed Effect

همچنین چون در این مدل از متغیرهای مجازی استفاده می‌شود و سپس روش OLS برای برآورد ضرایب آن به کار می‌رود؛ لذا آن را روش حداقل مربعات متغیرهای مجازی<sup>۱</sup> (LSDV) نیز می‌گویند. به این منظور مجموع مجذور خطاها را برای مدل می‌نویسیم.

$$RSS_{LSDV} = RSS_{FE} = \sum_i \sum_t e_{it}^2 = \sum_i \sum_t (Y_{it} - \hat{a}_i - \hat{B} X_{it})^2 \quad (11)$$

با مشتق‌گیری معادله (۱۱) نسبت  $a_i$  به  $B$  و خواهیم داشت:

$$\frac{\partial RSS_{FE}}{\partial \hat{a}_i} = -2 \sum_t (Y_{it} - \hat{a}_i - \hat{B} X_{it}) = 0 \quad (12)$$

$$\frac{\partial RSS_{FE}}{\partial \hat{B}} = -2 \sum_i \sum_t (Y_{it} - \hat{a}_i - \hat{B} X_{it}) X_{it} = 0 \quad (13)$$

با حل معادلات (۱۲) و (۱۳)،  $a_i$  و  $B$  به دست می‌آید.

$$\hat{a}_i = \overline{Y_{io}} - \hat{B} \overline{X_{io}} \quad (14)$$

$$\hat{B}_{LSDV} = \frac{\sum_i \sum_t (X_{it} - \overline{X_{io}})(Y_{it} - \overline{Y_{io}})}{\sum_i \sum_t (X_{it} - \overline{X_{io}})^2} \quad (15)$$

به طوری که:

$$\overline{Y_{io}} = \frac{\sum_t Y_{it}}{T} \quad \overline{X_{io}} = \frac{\sum_t X_{it}}{T} \quad (16)$$

را میانگین گروهی می‌گویند. ملاحظه می‌شود که نتایج فوق مانند مدل تجمیعی است با این تفاوت که به جای میانگین‌های کل از میانگین‌های گروهی استفاده شده است [۴].

**مدل اثرات تصادفی<sup>۲</sup>** مدل اثرات ثابت، امکان بررسی فردی مشاهده‌نشده را که با متغیرهای توضیحی همبستگی دارند فراهم می‌کند. در چنین شرایطی می‌توان تفاوت‌های فردی را به عنوان انتقال تابع رگرسیون تصور نمود. این مدل عمدتاً برای بررسی خصوصیات فردی یا گروهی واحدهای مورد مطالعه قابل کاربرد است؛ بنابراین نمی‌توان نتایج آن را به واحدهای خارج از نمونه تعمیم داد؛ زیرا اثرات ثابت مختص هر فرد یا گروه است که سایر افراد یا گروه‌ها، فاقد آن هستند. به همین دلیل است که آن را اثرات ثابت می‌نامند؛ یعنی خصوصیات فردی در طول زمان تغییر نمی‌کند. اگر اثرات فردی یا گروهی اکیدا با متغیرهای توضیحی همبستگی نداشته باشد در این صورت باید جملات ثابت فردی ( $a_i$ ) را به نحوی مدل‌سازی نمود تا به صورت تصادفی در بین گروه‌ها توزیع شود. در این جا خصوصیات فردی یا گروهی، ارتباطی با متغیرهای توضیحی ندارند؛ زیرا تصادفی هستند. اگر بر این باور باشیم که گروه‌ها از یک جامعه بزرگ نمونه‌گیری شده‌اند به نظر می‌آید که این روش، مناسب

<sup>1</sup> Least Squares Dummy Variables

<sup>2</sup> Random Effect

است. یکی از نتایج مهم این روش، آن است که تعداد ضرایب را تقلیل می‌دهد؛ زیرا در مدل اثرات ثابت باید  $n$  متغیر مجازی تعریف کنیم که درجه آزادی را کاهش می‌دهد؛ بنابراین مدل رگرسیون آن به صورت معادله (۱۷) می‌شود:

$$Y_{it} = \sum_{k=1}^K B_k X_{KIT} + (a + u_i) + \varepsilon_{it} \quad (17)$$

در این مدل  $K$  متغیر توضیحی، به علاوه یک جمله ثابت  $a$  داریم. در این مدل جمله ثابت  $a$ ، بیانگر میانگین ناهمگنی‌ها یا تفاوت‌های مشاهده نشده است. ناهمگنی‌های فردی یا گروهی را با  $z_i'a$  و متوسط آن را با  $E(z_i'a)$  نشان می‌دهیم؛ بنابراین  $u_i = z_i'a - E(z_i'a)$  است. در واقع  $u_i$  شامل مجموعه عواملی است (یعنی  $z_i'a$ ) که در رگرسیون نیستند؛ ولی مختص هر گروه می‌باشند [۴].

### ۳-۲ روش شبکه عصبی مصنوعی چند لایه پرسپترون

شبکه عصبی مصنوعی سعی در مدل کردن درک انسان، از طریق شبیه‌سازی فرآیند یادگیری دارد که بر اساس درک انسان پایه‌گذاری شده است. شبکه‌های عصبی مصنوعی هوشمند برای پیش‌بینی ابزار مناسبی هستند. به خصوص برای پیش‌بینی اطلاعاتی که ممکن است غیرساختاری نیز باشند. یکی از کاربردهای سیستم‌های عصبی مصنوعی، می‌تواند پیش‌بینی تقاضای وجه نقد دستگاه‌های خودپرداز باشد. مادامی که این کار به صورت فعلی انجام می‌شود هزینه‌های نگهداری، حمل و تکمیل موجودی دستگاه‌های خودپرداز بسیار پر هزینه خواهد بود که جنبه مهمی از سودآوری بانک را تحت تاثیر قرار می‌دهد. سیستم عصبی مصنوعی باید اطلاعات مربوط به میزان برداشت وجه نقد دستگاه‌های خودپرداز را که به صورت پنل هستند به عنوان ورودی آموزش دهد و همچنین باید نیازهای واقعی میزان تقاضای وجه نقد دستگاه‌های خودپرداز را به عنوان ستانده مطلوب بیاموزد. فرآیند کلی به کارگیری شبکه عصبی در حل مساله پیش‌بینی تقاضای وجه نقد دستگاه‌های خودپرداز، شامل گام‌های زیر است:

- جمع‌آوری همه اطلاعات در یک جا: برای توسعه مدل شبکه عصبی مصنوعی باید اطلاعات مربوط به دستگاه‌های خودپرداز به صورت یک پایگاه داده جمع‌آوری گردد. خلاصه‌ای از اطلاعات مربوط به دستگاه‌های خودپرداز عبارتند از: اسم دستگاه خودپرداز، محل قرارگیری دستگاه خودپرداز، ظرفیت پول و میزان برداشت پول در روز.

- تفکیک اطلاعات به مجموعه‌های آموزشی و آزمایشی: در حالت کلی برای شبکه عصبی در ابتدا یک مجموعه برای آزمایش استخراج می‌گردد، آنگاه یک مجموعه آموزش از نمونه‌های باقی‌مانده انتخاب می‌شود که در مدل ما ۶۰٪ داده‌ها برای آموزش و ۲۰٪ برای آزمایش و ۲۰٪ برای اعتبارسنجی استفاده شده است.
- تبدیل داده‌ها به ورودی‌های مناسب برای شبکه عصبی مصنوعی: پایگاه داده گردآوری شده، شامل دو نوع داده است که عبارتند از: الف) داده‌های عددی مثل ظرفیت پول، میزان برداشت پول در روز و ... ب) داده‌های کاراکتری مثل اسم دستگاه خودپرداز، محل قرارگیری دستگاه خودپرداز و ... . اطلاعات مورد نیاز برای شبکه عصبی مصنوعی، میزان برداشت وجه نقد در طی روزهای سال است که به عنوان ورودی برای شبکه در نظر

گرفته می‌شود و اطلاعاتی مانند اسم و آدرس محل قرارگیری می‌تواند برای جمع‌آوری اطلاعات اضافی مفید باشد؛ ولی نمی‌تواند به عنوان ورودی برای شبکه، مورد استفاده قرار گیرد.

آموزش و آزمایش شبکه عصبی مصنوعی: شبکه عصبی مصنوعی چند لایه با آموزش پس انتشار خطا برای پیش‌بینی در نظر گرفته می‌شود. انتخاب یک معماری مناسب برای شبکه عصبی یکی از حوزه‌هایی است که تحقیقات زیادی در مورد آن صورت گرفته است [۱۷]. هدف از یادگیری پس انتشار به دست آوردن یک الگوریتم آموزش مناسب است که وزن‌های سیناپسی شبکه را محاسبه نماید به گونه‌ای که تابع هزینه مینیم شود [۱۸]. در ابتدا فرض می‌کنیم که شبکه شامل یک تعداد ثابت از  $L$  لایه نرون است.  $K_0$  نرون در لایه ورودی و  $K_r$  نرون در  $r$  امین لایه، به طوری که  $r = 1, 2, \dots, L$  باشد. در همه نرون‌ها تابع فعالیت زیگمویید به کار برده شده است.  $N$  جفت آموزشی  $(y(i), x(i))$  در دسترس هستند به طوری که  $i = 1, 2, \dots, N$  می‌باشد چون ما فرض کرده‌ایم که  $K_L$  تعداد نرون‌های خروجی است؛ بنابراین خروجی یک عدد اسکالر اما به صورت یک بردار  $K_L$  بعدی خواهد بود.

$$y(i) = [y_{1(i)}, \dots, y_{K_L(i)}]^T$$

بردارهای ورودی به صورت بردارهای  $K_0$  بعدی هستند.

$$x(i) = [x_{1(i)}, \dots, x_{K_0(i)}]^T$$

در طول آموزش وقتی برابر  $x(i)$  به عنوان ورودی به کار برده می‌شود خروجی شبکه به صورت  $\hat{y}(i)$  خواهد شد که متفاوت از مقدار مطلوب  $y(i)$  است. وزن‌های سیناپسی محاسبه می‌شود به طوری که تابع هزینه  $J$  که به مقادیر  $y(i)$  و  $\hat{y}(i)$  بستگی دارد مینیمم شود؛ بنابراین مینیمم‌سازی تابع هزینه می‌تواند از طریق تکنیک‌های مناسب انجام شود. در این جا ما از برنامه کاهش گرادیان استفاده خواهیم کرد. فرض کنید که  $W_j^r$  بردارهای وزن  $j$  امین نرون در  $r$  امین لایه باشد که یک بردار  $K_{r-1} + 1$  بعدی است و به این صورت تعریف می‌شود:

$$W_j^r = [W_{j,1}^r, W_{j,2}^r, \dots, W_{j,K_{r-1}}^r]^T$$

فرمول کلی آن به صورت معادله (۱۸) و (۱۹) خواهد بود:

$$W_j^r(\text{new}) = W_j^r(\text{old}) + \Delta W_j^r \quad (18)$$

$$\Delta W_j^r = -\mu \frac{\partial J}{\partial W_j^r} \quad (19)$$

به طوری که  $w_j^r(\text{old})$  تخمین جاری از وزن ناشناخته است و  $\Delta W_j^r$  تصحیح، برای به دست آوردن تخمین بعدی  $w_j^r(\text{new})$  است. تابع هزینه به فرم معادله (۲۰) است:

$$J = \sum_{i=1}^N \mathcal{E}(i) \quad (20)$$

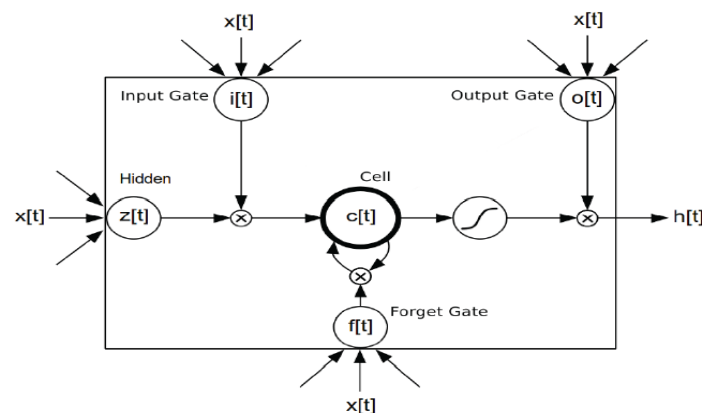
به طوری که  $\mathcal{E} \in$  یک تابع تعریف شده است که به  $\hat{y}(i)$  و  $y(i)$  بستگی دارد به طوری که  $i = 1, 2, \dots, N$  است. به عبارت دیگر  $\mathcal{E}$  به صورت یک جمع از  $N$  مقدار بیان می‌شود که تابع  $\mathcal{E}$  برای هر کدام از جفت‌های آموزشی  $(y(i), x(i))$  محاسبه می‌شود. برای مثال ما می‌توانیم  $\mathcal{E}$  را به گونه‌ای انتخاب کنیم که مجموعی از خطاهای مربع شده در نرون‌های خروجی باشد.

$$\varepsilon(i) = \frac{1}{2} \sum_{m=1}^{k_i} e_m^2(i) = \frac{1}{2} \sum_{m=1}^{k_i} (y_m(i) - \hat{y}_m(i))^2 \quad (21)$$

برای محاسبه جمله تصحیحی در رابطه با محاسبه وزن جدید، گرادیان تابع هزینه  $J$  نسبت به وزن‌ها و به دنبال آن ارزیابی  $\frac{\partial \varepsilon(i)}{\partial w_j^r}$  نیاز است [۱۹].

### ۳-۳ روش شبکه عصبی مصنوعی بازگشتی عمیق LSTM

شبکه‌های عصبی LSTM از لایه‌هایی از نرون‌ها تشکیل شده‌اند. داده ورودی از طریق شبکه برای ایجاد یک پیش‌بینی انتشار داده می‌شود. همانند شبکه‌های عصبی بازگشتی<sup>۱</sup>، شبکه‌های عصبی عمیق LSTM هم، دارای اتصالات بازگشتی هستند؛ بنابراین وضعیت فعالیت‌های قبلی نرون از گام زمانی قبلی به عنوان بخشی از داده‌ها برای فرموله کردن ورودی استفاده می‌شود؛ اما برخلاف شبکه‌های عصبی بازگشتی، شبکه عصبی LSTM یک ویژگی منحصربه‌فرد دارد که باعث می‌شود تا در آن، مشکلاتی که در آموزش و مقیاس‌بندی شبکه عصبی بازگشتی سنتی وجود دارد اجتناب شود. به همین دلیل و نتایج موثری که می‌تواند توسط شبکه عصبی بازگشتی عمیق LSTM ایجاد شود دلایلی هستند که باعث محبوبیت تکنیک شبکه عصبی مصنوعی LSTM است. چالشی که در شبکه عصبی بازگشتی سنتی با آن مواجه هستیم آن است که چگونه آن را به طور موثری آموزش دهیم. آزمایش‌ها نشان می‌دهد که شبکه عصبی بازگشتی دچار اضمحلال گرادیان می‌شود و همچنین به میزان زیادی ممکن است دچار مشکل بیش‌برازش شود که LSTM می‌تواند بر این مشکلات فائق آید. رنج اطلاعاتی که شبکه عصبی بازگشتی در عمل می‌تواند با آن مواجه شود محدود است. چون تاثیر یک ورودی داده شده روی لایه‌های مخفی و در نتیجه روی خروجی شبکه، ده‌ها برابر یا به صورت توانی در چرخه‌های بازگشتی شبکه افزایش می‌یابد. شبکه عصبی عمیق LSTM یک نوع معماری خاص از شبکه عصبی بازگشتی است که قادر به مواجه شدن با این مشکلات است. ساختار شبکه عصبی مصنوعی بازگشتی عمیق از نوع LSTM در شکل ۱ نشان داده شده است:



شکل ۱. ساختار LSTM با گیت فراموشی [۲۰]

<sup>1</sup> Recurrent Neural Network

در این تحقیق ما از LSTM با گیت‌های فراموشی که توسط فلیکس و جورگن و فرد معرفی شده‌اند [۲۰] استفاده کردیم. این نوع از شبکه‌ها دارای یک ساختار اضافی گیت فراموشی هستند. گیت فراموشی شبکه را قادر می‌سازد تا وضعیت خود را راه اندازی مجدد نماید.

$$\text{Block input unit : } z^t = g(W_z X^t + R_z Y^{t-1} + b_z) \quad (22)$$

در معادله (۲۲)،  $g$  زیگموید لوجستیک است.  $W$  ماتریس وزن از ورودی به  $Z$  می‌باشد و  $X$  گام ورودی جاری و  $R$  ماتریس وزن از خروجی گام قبلی به  $Z$  و  $b$  بردار بایاس است.

$$\text{Input gate unit : } i^t = \sigma(W_i X^t + R_i Y^{t-1} + b_i) \quad (23)$$

در معادله (۲۳)،  $\sigma$  زیگموید لوجستیک و  $W$  ماتریس وزن از ورودی به  $i$  و  $X$  گام ورودی جاری و  $R$  ماتریس وزن از خروجی گام قبلی به  $i$  و  $b$  بردار بایاس است.

$$\text{Forget gate unit : } f^t = \sigma(W_f X^t + R_f Y^{t-1} + b_f) \quad (24)$$

در معادله (۲۴)،  $\sigma$  زیگموید لوجستیک و  $W$  ماتریس وزن از ورودی به  $f$  و  $X$  گام ورودی جاری و  $R$  ماتریس وزن از خروجی گام قبلی به  $f$  و  $b$  بردار بایاس است.

$$\text{Cell state unit : } c^t = i^t \odot z^t + f^t \odot c^{t-1} \quad (25)$$

در معادله (۲۵)،  $i$  پارامترهای گیت ورودی و  $\odot$  ضرب نقطه‌ای دو بردار و  $f$  پارامترهای گیت فراموشی و  $c^{t-1}$  وضعیت سلول از گام قبلی است.

$$\text{Output gate unit : } o^t = \sigma(W_o X^t + R_o Y^{t-1} + b_o) \quad (26)$$

در معادله (۲۶)،  $\sigma$  زیگموید لوجستیک و  $W$  ماتریس وزن از ورودی به  $o$  و  $X$  گام ورودی جاری و  $R$  ماتریس وزن از خروجی گام قبلی به  $o$  و  $b$  بردار بایاس است.

$$\text{lock output unit : } Y^t = (o^t \odot h(c^t)) \quad (27)$$

در معادله (۲۷)،  $o$  پارامترهای گیت خروجی و  $\odot$  ضرب نقطه‌ای دو بردار و  $h$  زیگموید لوجستیک و  $c$  پارامتر وضعیت سلول است. شبکه عصبی بازگشتی LSTM دارای دو گذرگاه جلورو و پس‌انتشار است. پس‌انتشار محبوب‌ترین الگوریتمی است که به میزان زیادی در فیلد شبکه عصبی استفاده شده است. مرحله پس‌انتشار یک تکنیک برپایه گرادیان برای شبکه عصبی بازگشتی است. بر اساس روش پس‌انتشار، پیش‌بینی نتایج در زمان  $T$  می‌تواند صحیح‌تر باشد اگر آن بتواند از طریق آن چه که از سیستم در مراحل زمانی قبل‌تر عبور داده شده است محاسبه شود [۲۱]. بر اساس نظر کلائوس و همکارانش، پس‌انتشار برای LSTM می‌تواند در دو مرحله انجام شود. مرحله اول محاسبه کردن دلتا در بلوک LSTM و در مرحله دوم، دلتای وزن‌ها محاسبه می‌شود [۲۰].

از معایب مربوط به شبکه‌های عصبی عمیق، مشکل بیش‌برازش و زمان محاسباتی بالاست. شبکه‌های عصبی عمیق در معرض بیش‌برازش هستند چون تعداد لایه‌های اضافه شده، این امکان را ایجاد می‌کند که وابستگی‌های نایاب در داده‌های آموزشی مدل شود. روش‌های تنظیم، مثل روش هرس کردن که توسط ایواخنکو [۲۲] بیان شده یا کاهش وزن می‌تواند در طول آموزش به کار برده شود تا به از بین بردن بیش‌برازش کمک کند [۲۳]. روش

دیگری که برای تنظیم، اخیراً به کار برده می‌شود روش dropout است که در این روش تعدادی از واحدها به صورت تصادفی از لایه‌های مخفی حذف می‌شوند [۲۴]. این کار به از بین بردن وابستگی‌های نادر کمک می‌کند که می‌تواند در داده‌های آموزشی رخ دهد.

روشی که معمولاً برای آموزش این نوع شبکه‌ها استفاده می‌گردد روش پس‌انتشار با کاهش گرادیان است [۲۵] که این روش آموزش، حجم محاسبات بالایی دارد. چند پارامتر آموزشی وجود دارد که برای شبکه عصبی عمیق باید در نظر گرفته شود مثل سائز (تعداد لایه‌ها و تعداد نرون‌ها در هر لایه)، نرخ یادگیری و وزن‌های آغازین. جستجو در فضای پارامترها، برای یافتن پارامترهای بهینه ممکن است به دلیل هزینه زمانی و محاسباتی امکان‌پذیر نباشد. راه‌های متنوعی مثل دسته‌بندی کردن (محاسبه گرادیان روی چند نمونه آموزش در یک بار به جای محاسبه برای تک تک نمونه‌ها) باعث می‌شود که سرعت محاسبات بالا رود [۲۶].

#### ۴ بحث پیرامون نتایج

در این پژوهش ما از سه روش برای پیش‌بینی میزان تقاضای وجه نقد دستگاه‌های خودپرداز استفاده کردیم که عبارتند از:

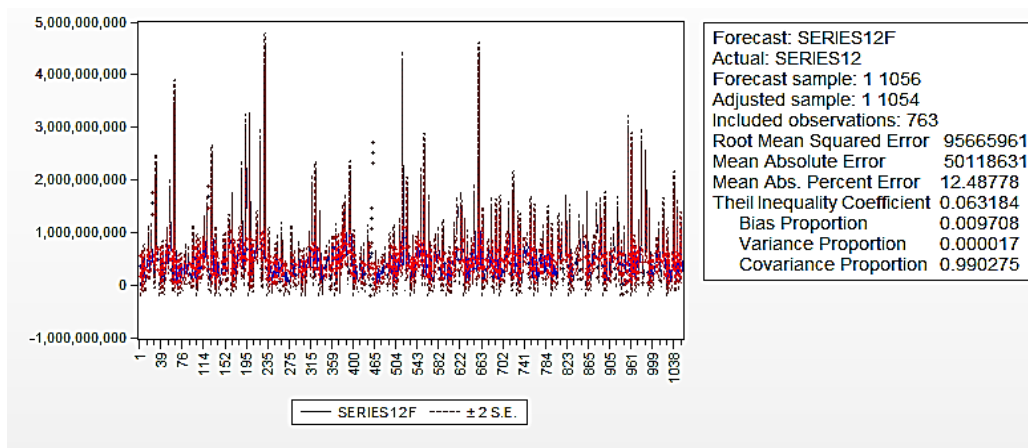
- روش سنتی آماری
- روش شبکه عصبی مصنوعی MLP
- روش شبکه عصبی بازگشتی عمیق LSTM

#### ۴-۱ نتایج روش سنتی آماری

روش‌های آماری متنوعی برای پیش‌بینی و تخمین مدل رگرسیونی وجود دارد که لازم است در ابتدا انواع داده‌ها شناسایی شود، سپس متناسب با نوع داده‌ها از مدل رگرسیونی مناسب استفاده شود. اگر داده‌های مساله مورد نظر؛ یعنی میزان برداشت وجه نقد دستگاه‌های خودپرداز را به صورت مقطعی در نظر بگیریم، در این صورت فقط به تفاوت‌های رفتار برداشت وجه نقد دستگاه‌های مختلف خودپرداز در یک روز توجه کرده‌ایم و تغییرات زمانی میزان برداشت وجه نقد در طول سال را نادیده گرفته‌ایم؛ بنابراین این نوع مدل داده‌ای یعنی مقطعی، مدل مناسبی برای پیش‌بینی تقاضای وجه نقد دستگاه‌های خودپرداز شبکه بانکی نیست. حال اگر مدل داده‌ای سری زمانی را برای پیش‌بینی میزان تقاضای وجه نقد دستگاه‌های خودپرداز در نظر بگیریم، به این ترتیب فرض کرده‌ایم که رفتار دستگاه‌های خودپرداز همگن هستند و بر روی هم هیچ تاثیری ندارند و فقط به تغییرات زمانی میزان برداشت وجه نقد هر دستگاه خودپرداز به تنهایی توجه کرده‌ایم؛ بنابراین بهترین مدل داده‌ای برای مساله مورد نظر، مدل داده‌ای ترکیبی است. مدل ترکیبی نیز خود به دو دسته تجمیعی و پنل قابل تفکیک است.

با توجه به نتایج آزمون آماری  $F$  که بر روی داده‌ها انجام شد و احتمال آزمون  $F$  که مقدار صفر را نشان می‌دهد، بیانگر این است که فرضیه  $H_0$  آزمون آماری  $F$  رد می‌شود و مدل رگرسیونی مناسب برای داده‌های موردنظر، مدل پنل است و مدل تجمیعی، روش رگرسیون کاملی برای بیان اثرات داده‌های میزان برداشت وجه

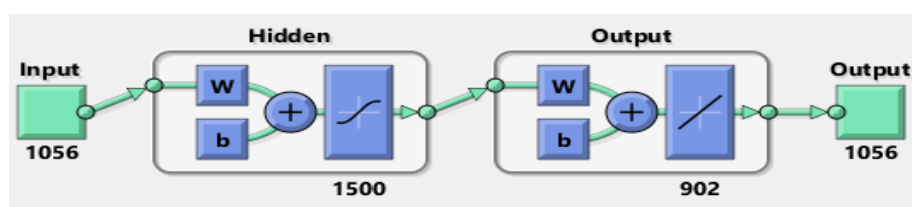
نقد دستگاه‌های خودپرداز شبکه بانکی نیست. با بررسی آزمون هاسمن بر روی داده‌های میزان برداشت وجه نقد دستگاه‌های خودپرداز، ما با اثرات ثابت در داده‌های پنل مواجه هستیم. ضرایب مدل رگرسیونی توسط نرم‌افزار Eviews برای داده‌های میزان برداشت وجه نقد دستگاه‌های خودپرداز تخمین زده شد به طوری که نتایج پیش‌بینی با استفاده از روش سنتی آماری (مدل رگرسیونی داده‌های پنل با اثرات ثابت) در شکل ۲ نشان داده شده است. با توجه به نتایج به دست آمده از پیش‌بینی، درصد خطا  $12/48778\%$  تخمین زده شده است.



شکل ۲. نتایج پیش‌بینی با استفاده از مدل آماری

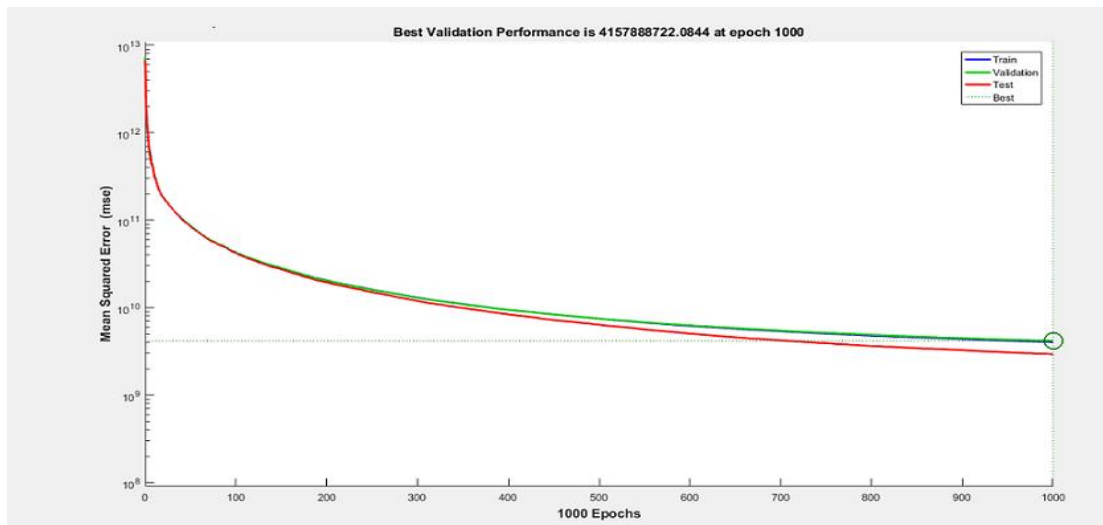
#### ۴-۲ روش شبکه عصبی مصنوعی MLP

ما در این تحقیق از نرم‌افزار متلب برای طراحی شبکه عصبی مصنوعی چند لایه استفاده کردیم. اطلاعات ورودی و خروجی مطلوب را، در ابتدا برای نرم افزار متلب تعریف کردیم. اطلاعات ورودی، داده‌های پنل مربوط به دستگاه‌های خودپرداز شبکه بانکی در سال ۱۳۹۳ و اطلاعات خروجی میزان تقاضای وجه نقد مطلوب مربوط به این دستگاه‌های خودپرداز است که این اطلاعات، همان میزان برداشت وجه نقد در سال ۱۳۹۴ است که به عنوان ستانده مطلوب به شبکه داده می‌شود. با استفاده از این اطلاعات، به آموزش و آزمایش شبکه عصبی مصنوعی چند لایه با کمک نرم افزار متلب پرداختیم. شبکه طراحی شده، در شکل ۳ نشان داده شده است و نتایج به دست آمده از این طراحی در شکل ۴ و ۵ نمایش داده شده است.

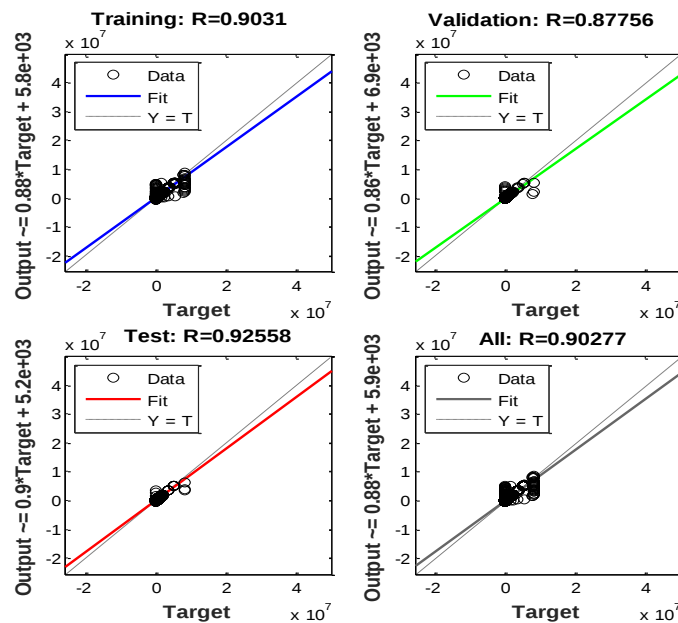


شکل ۳. شبکه عصبی مصنوعی طراحی شده به کمک نرم افزار متلب

پور ذاکر عربانی و ابراهیم پور کومله، بهینه‌سازی پیش‌بینی تقاضای وجه نقد دستگاه‌های خودپرداز شبکه بانکی...



شکل ۴. نمودار خطای MSE برای داده‌های آموزشی و آزمایشی



شکل ۵. نمودار رگرسیون با صحت ۹۲/۵۵۸

با توجه به نتایج به‌دست آمده از آزمایش شبکه عصبی مصنوعی با داده‌های پنل میزان برداشت دستگاه‌های خودپرداز شبکه بانکی با استفاده از نرم افزار متلب، ما به صحت ۹۲/۵۵۸٪ برای پیش‌بینی تقاضای وجه نقد دستگاه‌های خودپرداز دست یافتیم. با بررسی نتایج به‌دست آمده از دو روش آماری و شبکه عصبی مصنوعی چند لایه، که در شکل ۲ و ۵ نشان داده شده است، مشخص می‌شود که روش هوشمند شبکه عصبی مصنوعی چند لایه عملکرد بهتری را برای پیش‌بینی مقادیر آتی میزان تقاضای وجه نقد دستگاه‌های خودپرداز نسبت به روش سنتی آماری از خود نشان می‌دهد.

#### ۴-۳ روش شبکه عصبی بازگشتی LSTM

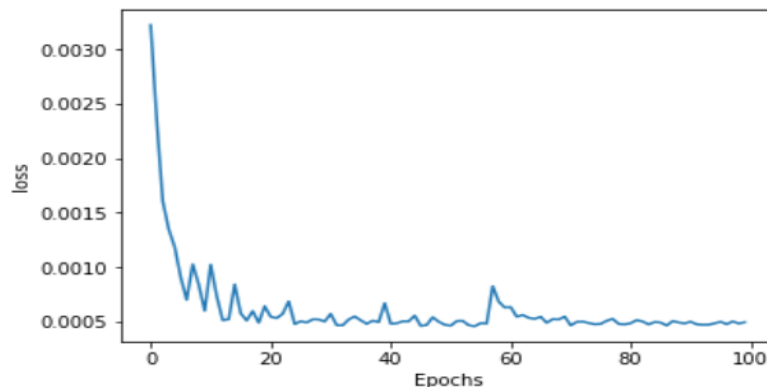
در مدل پیشنهادی برای شبکه عصبی عمیق LSTM، ما به کشف وابستگی‌های زمانی برای ۲ سال برای آموزش شبکه، نیاز داریم. دلیل انتخاب ۲ سال برای آموزش شبکه این است که بتوانیم اثر روزهای خاص و تعطیلات، در اختیار شبکه قرار داده شود. به دلیل متفاوت بودن تعداد روزهای سال شمسی و قمری و از طرفی چون بیش‌تر اعیاد و تعطیلات رسمی کشور ما مذهبی هستند و این تعطیلات از نظر تقویمی تغییراتی ۱۰ روزه دارند، با در نظر گرفتن میزان برداشت‌های ۲ سال برای دستگاه‌های خودپرداز، این تغییرات در اطلاعاتی وجود دارد که سیستم با آن آموزش داده می‌شود و از طرفی چون داده‌های ما از نوع پنل هستند؛ یعنی بین دستگاه‌های خودپرداز، تعاملات و تاثیرگذاری‌هایی از نظر نزدیکی جغرافیایی وجود دارد به همین خاطر با استفاده از شبکه عصبی بازگشتی LSTM، شبکه عصبی عمیقی را طراحی کردیم که قدرت تشخیص ویژگی‌های مورد نظر؛ یعنی وابستگی‌های زمانی تا چند صد مرحله را داشته باشد و مدل مناسبی برای داده‌های پنل با حجم زیاد باشد.

در مدل پیشنهادی این تحقیق، LSTM را از طریق مقایسه معماری‌های مختلف (یعنی تعداد لایه‌های مخفی و تعداد نرون‌ها در هر لایه)، از نظر صحت خروجی‌ها بهینه‌سازی کردیم. چون بهینه‌سازی تعداد لایه‌ها و تعداد نرون‌ها در هر لایه مخفی برای عملیات مختلف هنوز یک مساله حل نشده در حوزه پژوهش شبکه عصبی مصنوعی می‌باشد. تعداد کم واحدهای مخفی باعث ایجاد خطای بیش‌تری می‌شود و شبکه همگرا نمی‌شود و از طرفی دیگر تعداد واحدهای مخفی زیاد نیز منجر به خطای زیاد می‌شود چون باعث ایجاد بیش‌تر برازش می‌شود. به همین خاطر تعداد بهینه نرون‌ها را از طریق سعی و خطا به دست آوردیم. شبکه عصبی بازگشتی عمیق LSTM را توسط نرم‌افزار پایتون همراه با توابع کتابخانه‌ای کراس به دو روش تفکیک داده‌های آزمایش و آموزش پیاده‌سازی کردیم. در ابتدا برای تفکیک داده‌ها به این صورت عمل کردیم که ۲۰٪ داده‌ها را برای آزمایش و ۸۰٪ داده را برای آموزش انتخاب کردیم. سپس ساختار و معماری شبکه عمیق را پیاده‌سازی کردیم که خلاصه‌ای از ساختار معماری شبکه در جدول ۱ نشان داده شده است. در این شبکه از سه لایه نرون استفاده شده به طوری که لایه اول و دوم از نوع LSTM هستند و در هر لایه از ۸۴۴ تا نرون استفاده شده است و لایه سوم از نوع اتصالات کامل است که در آن از ۳۶۶ نرون استفاده شده است.

جدول ۱. ساختار شبکه LSTM

| Layer (type)                 | Output Shape | Param # |
|------------------------------|--------------|---------|
| Lstm_1 (LSTM)                | (4,1,844)    | 4088336 |
| Lstm_2 (LSTM)                | (4,1,844)    | 5702064 |
| Dense_1 (Dense)              | (4,1,366)    | 309270  |
| Total params: 10,099,670     |              |         |
| Trainable params: 10,099,670 |              |         |
| Non-trainable params:0       |              |         |

سپس به آموزش شبکه طراحی شده پرداختیم و به تعداد ۱۰۰ مرحله آموزش را تکرار کردیم. نتایج حاصل از آن در شکل ۶ نشان داده شده است.



شکل ۶. منحنی تابع Loss شبکه عصبی عمیق بازگشتی LSTM در ۱۰۰ Epoch

صحت به دست آمده از این مدل برای داده‌های آموزش ۹۹/۶۵٪ و برای داده‌های تست ۹۶/۶۵٪ شد. در مرحله بعد برای تفکیک داده‌ها، از روش cross validation با  $kfold=5$  استفاده کردیم. سپس به آموزش و آزمایش شبکه پرداختیم. صحت به دست آمده از این مدل برای داده‌های آموزشی ۱۰۰٪ و برای داده‌های آزمایشی ۹۸/۸۶٪ شد.

## ۵ نتیجه‌گیری و پیشنهادها

یکی از مشکلات مدیران بانکی، دقت در پیش‌بینی میزان وجه نقد مورد نیاز برای دستگاه‌های خودپرداز است. میزان تقاضای وجه نقد دستگاه‌های خودپرداز به میزان برداشت‌های قبلی وجه نقد مرتبط است. دقت و صحت در پیش‌بینی میزان تقاضای مورد نیاز وجه نقد در دستگاه‌های خودپرداز به دو دلیل حایز اهمیت است. اگر میزان پول موجود در دستگاه بیش‌تر از حد نیاز باشد بخش قابل توجهی از پول بدون استفاده باقی می‌ماند و با توجه به تورم قابل ملاحظه‌ای که در کشور وجود دارد بی‌استفاده ماندن پول در دستگاه، هزینه‌هایی از جهت بی‌ارزش شدن پول خواهد داشت و از طرف دیگر اگر پول نقد به اندازه کافی در دستگاه خودپرداز نباشد باعث نارضایتی مشتریان این سیستم بانکی می‌شود که این وضعیت نیز برای سیستم بانکی به دلیل از دست دادن بخشی از مشتریان خود هزینه‌هایی را خواهد داشت. هدف از انجام این تحقیق دستیابی به اهداف علمی و کاربردی است که در ادامه به طور مختصر به آن اشاره می‌کنیم. هدف علمی ما از این پژوهش، آنالیز حجم عظیمی از داده‌های پنل موجود و پیش‌بینی مقادیر آینده آن‌ها در سیستم دستگاه خودپرداز در شبکه بانکی بود که برای دست یافتن به این منظور، به مدل کردن یک شبکه عصبی بازگشتی عمیق LSTM پرداختیم که هم قابلیت تشخیص ویژگی وابستگی‌های زمانی بیش از چند صد مرحله را داشته باشد و هم ارتباطات مقطعی بین دستگاه‌های خودپرداز را تشخیص دهد و سپس به بررسی مقایسه‌ای نتایج پیش‌بینی روش پیشنهادی این پژوهش با روش سنتی آماری و روش شبکه عصبی مصنوعی MLP پرداختیم و هدف کاربردی ما کاهش هزینه‌های دوگانه سیستم بانکی برای پیش‌بینی دقیق میزان تقاضای وجه نقد دستگاه‌های خودپرداز است.

در روش سنتی آماری در ابتدا به تشخیص نوع داده‌ها پرداختیم تا با تشخیص مناسب نوع داده‌ها بتوانیم مدل رگرسیونی مناسبی را برای پیش‌بینی تشخیص دهیم. برای همین منظور از دو نوع آزمون آماری (آزمون F و آزمون هاسمن) استفاده کردیم. در ابتدا بر روی داده‌ها آزمون F را انجام دادیم آزمون F دارای دو فرضیه است. اگر احتمال آزمون F بزرگ‌تر از ۰.۵٪ شود فرض صفر رد نمی‌شود و باید از مدل رگرسیونی مربوط به داده‌های تجمیعی برای پیش‌بینی استفاده نماییم. در غیر این صورت مدل داده‌های ما از نوع پنل خواهد بود. با توجه به نتایج به‌دست آمده از این آزمون بر روی داده‌های مربوط به میزان برداشت وجه نقد از دستگاه‌های خودپرداز که صفر درصد شد، مشخص شد که نوع داده‌های ما از نوع پنل هستند. پس از آن از آزمون هاسمن استفاده کردیم چون داده‌های پنل نیز بر دو نوع هستند یا دارای اثرات ثابت هستند و یا دارای اثرات تصادفی. آزمون هاسمن نیز دارای دو فرضیه است. اگر احتمال آزمون هاسمن کوچک‌تر از ۰.۵٪ شود، فرض صفر رد می‌شود و باید از مدل رگرسیونی داده‌های پنل با اثرات ثابت استفاده کرد. در غیر این صورت مدل رگرسیونی با اثرات تصادفی برای پیش‌بینی مناسب است. با توجه به آزمون‌های انجام شده، نوع داده‌های ما پنل، با اثرات ثابت تشخیص داده شد؛ بنابراین از مدل رگرسیونی پنل با اثرات ثابت برای پیش‌بینی استفاده شد. نتایج پیش‌بینی با استفاده از این مدل رگرسیونی میزان خطای ۱۲/۴۸۷۷۸٪ را نشان داد. روش بعدی انجام شده برای پیش‌بینی میزان تقاضای وجه نقد دستگاه‌های خودپرداز، روش شبکه عصبی مصنوعی پرسپترون چند لایه با استفاده از روش آموزش پس‌انتشار بود که نتایج به‌دست آمده از این روش، میزان صحت ۹۰/۳۱٪ برای داده‌های آموزشی و ۹۲/۵۵۸٪ را برای داده‌های آزمایش نشان داده است. روش سوم که روش پیشنهادی ما در این پژوهش است استفاده از شبکه عصبی بازگشتی عمیق LSTM است. نتایج به‌دست آمده از این روش با استفاده از روش تفکیک اطلاعات cross validation، میزان صحت ۱۰۰٪ را برای داده‌های آموزش و ۹۸/۸۶٪ را برای داده‌های آزمایش نشان داد. با بررسی مقایسه‌ای بین نتایج به‌دست آمده از روش‌های مختلف به این نتیجه رسیدیم که مدل پیشنهادی؛ یعنی شبکه عصبی بازگشتی عمیق LSTM، بهترین عملکرد را از نظر صحت و درستی پیش‌بینی از خود نشان داده است.

با توجه به نتایج بسیار مناسب روش شبکه عصبی مصنوعی بازگشتی LSTM بر روی داده‌های پنل با حجم زیاد، پیشنهاد می‌شود که از این روش که توانایی خوبی در مواجهه با داده‌هایی با وابستگی‌های زمانی (سری زمانی) و در ضمن تاثیرگذاری‌هایی از نوع مقطعی به صورت همزمان با وابستگی زمانی دارد، استفاده شود. در پیش‌بینی مسایل اقتصادی کشور از جمله پیش‌بینی قیمت سهام شرکت‌ها در بورس، پیش‌بینی قیمت طلا، پیش‌بینی قیمت نفت و ... نیز می‌توان از روش شبکه عصبی بازگشتی عمیق LSTM استفاده کرد.

## منابع

- [۱۴] عباسی ابراهیم، رستگاریا فاطمه، ابراهیمی فهیمه. (۱۳۹۳). پیش‌بینی نقدینگی موردنیاز دستگاه‌های خودپرداز با استفاده از مدل خطی (ARIMA) و غیرخطی (شبکه‌های عصبی). فصلنامه سیاست‌های راهبردی و کلان، ۲ (۸)، ۵۹-۷۶.
- [۱۵] تقوی فرد محمدتقی، خاتمی فیروزآبادی سیدمحمدعلی، سجادی سیدخلیل اله، رمضانین بادی اعظم. (۱۳۹۵). افزایش میزان رضایت شهروندان از دستگاه‌های خودپرداز بانک شهر و کاهش هزینه‌های اقتصادی بانک با به کارگیری مدل کنترل موجودی شبیه سازی شده. اقتصاد و مدیریت شهری ۵ (۴)، ۱-۱۸.

- [۱۶] حاجی مولانا، معمار پور، سجادی، سید خلیل‌الله. (۱۳۹۶). طراحی مدل پویای پیش‌بینی تقاضای پول در دستگاه‌های خودپرداز شهر تهران (مطالعه موردی: بانک شهر). نشریه مهندسی صنایع. ۲۳ (۵۱)، ۲۸۲-۲۹۵.
- [۱۸] فلاح جلودار مهدی. (۱۳۹۵). ارزیابی کارایی شرکت های توزیع نیروی برق ایران با استفاده از مدل ترکیبی شبکه‌های عصبی و تحلیل پوششی داده‌ها. تحقیق در عملیات در کاربردهای آن. ۱۳ (۴)، ۶۷-۸۳.
- [۱۹] سعید محرابیان، صابر ساعتی مهدی، علی هادی. (۱۳۹۰). ارزیابی کارایی شعب بانک اقتصاد نوین با ترکیبی از شبکه عصبی و تحلیل پوششی داده‌ها. تحقیق در عملیات در کاربردهای آن. ۸ (۴)، ۲۹-۳۹.
- [1] Siami Namini S. and Siami Namini A., (2018). Forecasting Economics and Financial Time Series: ARIMA vs. LSTM, arXiv:1803.06386.
- [2] Osorio A. and Toro H., (2012). An MIP model to optimize a Colombian cash supply chain. International Transactions in Operational Research, 19, 659-673.
- [3] Castro J. (2009). A stochastic programming approach to cash management in banking. European Journal of Operational Research, 192, 963-974.
- [4] Brooks C. Introductory econometrics for finance. (2014).
- [5] Teddy S. and Ng S., (2011). Forecasting ATM cash demands using a local learning model of cerebellar associative memory network. International Journal of Forecasting, 27, 760-776.
- [6] Andrawis R. R., Atiya A. F. and El-Shishiny H., (2011). Combination of long term and short term forecasts, with application to tourism demand forecasting. International Journal of Forecasting, 27, 870-886.
- [7] Venkatesh K., Ravi V., Prinzie A. and Van den Poel D., (2014). Cash demand forecasting in ATMs by clustering and neural networks. European Journal of Operational Research, 232, 383-392.
- [8] Simutis R., Dilijonas D., Bastina L., Friman J. and Drobinov P., (2007). Optimization of cash management for ATM network. Information technology and control. 36.
- [9] Broda P., Levajković T., Kresoja M., Marčeta M., Mena H., Nikolić M. and Stojančević T., (2014). Optimization of ATM filling-in with cash.
- [10] Dandekar P. V. and Ranade K. M., (2015). ATM Cash Flow Management. International Journal of Innovation, Management and Technology, 6, 343.
- [11] Garcia-Pedrero A. and Gomez-Gil P., (2010). Time series forecasting using recurrent neural networks and wavelet reconstructed signals. Electronics, Communications and Computer (CONIELECOMP), (2010). 20th International Conference on. 169-173.
- [12] Catal C., Fenerci A., Ozdemir B. and Gulmez O., (2015). Improvement of Demand Forecasting Models with Special Days. Procedia Computer Science, 59, 262-267.
- [13] Jadwal P. K., Jain S., Gupta U. and Khanna P., (2017). K-Means Clustering with Neural Networks for ATM Cash Repository Prediction. International Conference on Information and Communication Technology for Intelligent Systems, 588-596.
- [17] Theodoridis S. and Koutroumbas K., (1999). Pattern recognition.
- [20] Greff K., Srivastava R. K., Koutník J., Steunebrink B. R. and Schmidhuber J., (2017). LSTM: A search space odyssey. IEEE transactions on neural networks and learning systems.
- [21] Van Der Voort M., Dougherty M. and Watson S., (1996). Combining Kohonen maps with ARIMA time series models to forecast traffic flow. Transportation Research Part C: Emerging Technologies, 4, 307-318.
- [22] Ivakhnenko A. G., (1971). Polynomial theory of complex systems. IEEE transactions on Systems, Man and Cybernetics, 1, 364-378.
- [23] Bengio Y., Boulanger-Lewandowski N. and Pascanu R., (2013). Advances in optimizing recurrent networks. Acoustics, Speech and Signal Processing (ICASSP). IEEE International Conference on, 8624-8628.
- [24] Srivastava N., Hinton G. E., Krizhevsky A., Sutskever I. and Salakhutdinov R., (2014). Dropout: a simple way to prevent neural networks from overfitting. Journal of machine learning research, 15, 1929-1958.
- [25] Schmidhuber J. (2015). Deep learning in neural networks: An overview. Neural network, 61, 85-117.
- [26] Liu W., Wang Z., Liu X., Zeng N., Liu Y. and Alsaadi F. E., (2017). A survey of deep neural network architectures and their applications. Neurocomputing, 234, 11-26.