

A Strategic Inventory Management Model Based on Deep Learning

Hamidreza Hajali¹, Mohammad Ali Afshar Kazemi², Adel Azar³, Abbas Toloui Ashlaghi⁴, Reza Radfar⁵

Received: 2024/11/06

Accepted ;2025/01/28

Extended Abstract

Introduction

Warehouse management, as a dynamic subset of any organization, can organize the needs of the supply and distribution chain within an intelligent framework. This article aims to address the modern challenges of industrial big data, particularly non-linear data from qualitative historical records, which are influential in inventory management. Despite their valuable content, these data types are often unexploited by classic methods like time series for forecasting and decision-making.

Literature Review

While traditional inventory models like EOQ and EPQ are well-established, they often fall short in handling the complexity and non-linearity of modern industrial big data. Artificial intelligence, especially deep learning, offers promising capabilities. However, a significant gap exists in studies that strategically integrate classic inventory principles with advanced deep learning architectures for enhanced predictive accuracy. This research bridges this gap by presenting a hybrid model.

1.PhD Student, Department of Industrial Management, SR.C., Islamic Azad University, Tehran, Iran.

2.Professor, Department of Industrial Management, CT.C., Islamic Azad University, Tehran, Iran (Corresponding Author) dr.mafshar@gmail.com

3.Professor, Department of Industrial Management, Tarbiat Modares University, Tehran, Iran.

4.Professor, Department of Industrial Management, SR.C., Islamic Azad University, Tehran, Iran.

5..Professor, Department of Industrial Management, SR.C., Islamic Azad University, Tehran, Iran.

How to cite this article: Hajali, H., Afshar Kazemi, M. A., Azar, A., Toloui Ashlaghi, A., & Radfar, R. (2025). A Strategic Inventory Management Model Based on Deep Learning. *Modern Management Engineering*, 11(1), 253-280. [In Persian]

 <https://doi.org/10.71652/JMEM.2025.1189687>

Research Methodology

This study proposes a strategic model based on the Long Short-Term Memory (LSTM) deep learning technique. This approach enables the use of non-linear industrial big data for order forecasting. To prepare the data for machine training and to process it more effectively according to inventory importance levels, we utilized the Economic Order Quantity (EOQ), Economic Production Quantity (EPQ), and ABC analysis methods. The LSTM network was then constructed and trained with the prepared data and specified parameters.

Results

After building and training the LSTM model, a satisfactory result with ۹۳% accuracy was achieved, aligning with the research objective. The findings indicate that integrating traditional approaches with deep learning enhances the performance of machine learning for complex big data analysis tasks. This synergy allows the model to capture intricate patterns that traditional methods alone would miss.

Discussion and Conclusion

This research demonstrates that a hybrid model combining classic inventory management techniques with deep learning provides a powerful strategic tool for modern organizations. The proposed LSTM-based model offers a practical and highly accurate solution for forecasting, enabling companies to better manage their inventory, reduce costs, and improve supply chain efficiency by leveraging the full potential of their valuable, yet complex, historical data.

Keywords: Inventory Management, Big Data, Artificial Intelligence, Deep Learning, Strategic Model.

JEL Classification: C45, M45, O32

مدل راهبردی مدیریت موجودی انبار مبتنی بر یادگیری عمیق

حمیدرضا حاجعلی^۱، محمدعلی افشار کاظمی^۲، عادل آذر^۳، عباس طلوعی اشلقی^۴، رضا رادفر^۵

تاریخ پذیرش: ۱۴۰۳/۱۲/۲۲

تاریخ دریافت: ۱۴۰۳/۰۸/۱۶

چکیده

هدف: ارائه یک مدل راهبردی برای مدیریت موجودی انبار با استفاده از یادگیری عمیق (LSTM) جهت بهره‌برداری از کلان‌داده‌های غیرخطی و کیفی صنعتی که توسط روش‌های کلاسیک قابل استفاده نیستند.

روش‌شناسی پژوهش: در این پژوهش، یک مدل یادگیری عمیق حافظه بلندمدت (LSTM) طراحی شد. داده‌ها با استفاده از روش‌های سنتی مقدار اقتصادی سفارش (EOQ)، مقدار اقتصادی تولید (EPQ) و تحلیل ABC برای آموزش شبکه آماده‌سازی گردید.

یافته‌ها: مدل LSTM پس از ساخت و آموزش، به دقت پیش‌بینی ۹۳ درصدی دست یافت که نشان‌دهنده موفقیت ادغام رویکردهای سنتی با هوش مصنوعی برای تحلیل کلان‌داده‌ها در مدیریت موجودی است.

اصالت / ارزش افزوده علمی: نوآوری این تحقیق در تلفیق روش‌های کلاسیک مدیریت موجودی با تکنیک‌های پیشرفته یادگیری عمیق است که راهکاری عملی برای حل چالش پیش‌بینی بر اساس داده‌های غیرخطی در زنجیره تأمین ارائه می‌دهد.

کلیدواژه‌ها: مدیریت موجودی، کلان داده، هوش مصنوعی، یادگیری عمیق، مدل راهبردی.

طبقه‌بندی موضوعی: C45, M11, O32

۱. دانشجوی دکتری، گروه مدیریت صنعتی، واحد علوم و تحقیقات، دانشگاه آزاد اسلامی، تهران، ایران.
۲. استاد، گروه مدیریت صنعتی، واحد تهران مرکزی، دانشگاه آزاد اسلامی، تهران، ایران (نویسنده مسئول). پست الکترونیکی: dr.mafshar@gmail.com
۳. استاد، گروه مدیریت صنعتی، دانشگاه تربیت مدرس، تهران، ایران.
۴. استاد، گروه مدیریت صنعتی، واحد علوم و تحقیقات، دانشگاه آزاد اسلامی، تهران، ایران.
۵. استاد، گروه مدیریت صنعتی، واحد علوم و تحقیقات، دانشگاه آزاد اسلامی، تهران، ایران. نحوه استناد به مقاله: حاجعلی، ح، افشار کاظمی، م، ع، آذر، ع، طلوعی اشلقی، ع، و رادفر، ر. (۱۴۰۴). مدل راهبردی مدیریت موجودی انبار مبتنی بر یادگیری عمیق. مهندسی مدیریت نوین، ۱۱(۱)، ۲۵۳-۲۸۰.

۱- مقدمه

مدیریت موجودی شامل برنامه‌ریزی، هماهنگی و کنترل فعالیت‌های مرتبط با گردش موجودی‌هاست (Toomey, 2000). نقش موجودی‌ها، نقش اصلی در موفقیت یا شکست زنجیره تأمین است. از این رو هماهنگی سطوح موجودی در سرتاسر زنجیره تأمین حائز اهمیت است (Aro-Gordon & Gupte, 2016). فضای کسب‌وکارها شاهد تحولات عمیق و اساسی شده است و بسیاری از صاحب‌نظران حوزه تجارت و اقتصاد بر این باورند که این تحولات حاکی از شروع دوره جدید یعنی انقلاب صنعتی چهارم است (Marr, 2016). چهارمین انقلاب صنعتی را با عنوان انقلاب دیجیتال نیز معرفی می‌کنند. ویژگی‌های این دوره فراگیر بودن اینترنت موبایل، حسگرهای کوچک‌تر، قوی‌تر و درعین حال ارزان‌تر، هوش مصنوعی و یادگیری ماشین است؛ مراکز عظیم دیتاهای غیرخطی حاصل گسترش آن‌ها است (Sanaei, 2016). لازم به ذکر است، در محیط رقابتی کسب‌وکارهای به وجود آمده در این حوزه، شرکت‌ها در مواجهه با مسئله کلان داده‌ها با چالش‌هایی مانند تصمیم‌گیری سریع برای چابکی سازمان خود و البته بهبود بهره‌وری روبه‌رو هستند، زیرا بسیاری از سیستم‌های تولید انبوه کلاسیک، توانایی و چابکی لازم جهت مدیریت و پاسخگویی به نیازهای ترند شده مشتریان خود را ندارند (Lee et al., 2014). همچنین، پیشرفت دانش و فناوری باعث گسترش بی‌سابقه کسب‌وکارهای داده محور شده است. به‌طوری‌که هر سازمان در طول عمر خود با حجم زیادی از این داده‌ها روبه‌رو است، بنابراین، این نکته را باید در نظر داشت که حجم عظیم داده به یکی از مهم‌ترین منابع اطلاعاتی برای هردو طرف مصرف‌کنندگان و کسب‌وکارها تبدیل شده است (Duan et al., 2013).

در این راستا، مقاله پیش رو با پیشنهاد مدلی راهبردی؛ اثربخشی پیاده‌سازی تکنیک‌های یادگیری عمیق را در پیش‌بینی موجودی کالا مورد بررسی قرار می‌دهد. در ادامه نیز مسئله دیگری که عملکرد تجزیه و تحلیل کلان داده‌های غیرخطی و کیفی را در قالب آموزش به ماشین ارزیابی می‌کند. پیش‌تر نیز برای پردازش کلان داده، از یکی از زیرشاخه‌های مهم یادگیری ماشین که به‌عنوان یادگیری عمیق شناخته می‌شود استفاده گردیده است (Albayrak et al., 2023). برآیند امکانات جدیدی که هوش مصنوعی در اختیار ما جهت تحلیل کلان داده‌های غیرخطی قرار می‌دهد، موجب کشف مسیرهای استراتژیک تجاری

برای سازمان‌ها می‌گردد. بدیهی است اشتباه در تعیین سطح موجودی کالا به سرعت منجر به از دست دادن فروش و یا افزایش هزینه نگهداری می‌گردد. مدیران ارشد مالی و اجرایی به اهمیت مدیریت موجودی کالا پی برده‌اند. اغلب مدیران اجرایی مطلع‌اند، به دست آوردن مقدار مناسب موجودی انبار ضروری است، زیرا نه تنها هزینه‌ها را کنترل می‌نماید بلکه معیاری برای تعیین سلامت کل سازمان است. تأمین‌کنندگان برتر قادرند سطح موجودی کالا را با بهره‌برداری از ابزارهای تحلیلی بین ۲۰ تا ۵۰ درصد بهبود دهند و در نتیجه صرفه‌جویی فراوانی داشته باشند (Fischer & Krauss, 2018)

از آنجاکه کنترل موجودی هوشمند، بسیاری از قطعات و تجهیزاتی که در انبارها بلااستفاده مانده‌اند را وارد چرخه تولید و یا فروش می‌نماید و هزینه‌های نگهداری را کاهش و نرخ سرمایه در گردش را افزایش می‌یابد، از این‌رو بدیهی است، پیش‌بینی نقطه بهینه میزان موجودی کالا توان سازمان را در جهت نیل به اهداف راهبردی خود بیشتر می‌سازد. تأمین به موقع، به اندازه و با قیمت مناسب کالاهای مورد نیاز سازمان، تعیین میزان خرید مقرون به صرفه و نگهداری مقدار بهینه از موجودی، یکپارچه کردن اطلاعات موجودی تمام انبارهای سازمان. کنترل اصالت کالاها از زمان شروع به تولید توسط تأمین‌کننده تا تحویل به انبار، هوشمندسازی اعلام نیازها و قراردادهای و ثبت سفارش‌ها، ردیابی کالاهایی که از انبار خارج شده تا آخرین مرحله از طول عمر مصرف کالا. توزیع به موقع کالاها با مقدار و مشخصات معین بین درخواست‌کننده، کنترل گردش کالاها به منظور تضمین تاریخ انقضای مصرف آن‌ها، کاهش اقلام دورریز برگشتی از نقاط مصرف از طریق هوشمندسازی، بررسی کارشناسی و تخصیص صحیح اقلام به مراکز تعمیر درون یا برون سیستمی، کاهش هزینه‌ها از طریق کم کردن سطح موجودی انبارها و اصلاح جانمایی انبارها و همچنین افزایش نرخ سرمایه در گردش با توجه به خارج نمودن موجودی راکد از تسهیلات همه از فواید به کارگیری تکنیک‌های هوش مصنوعی در کنترل موجودی می‌باشند.

پیش‌بینی نیازهای پنهان ذینفعان در واقع در ادامه فرآیند یادگیری عمیق توسط خود ماشین صورت خواهد گرفت. انتظار می‌رود یادگیری عمیق در گامی فراتر با در اختیار گرفتن حجم عظیم اطلاعات دقیق بتواند چه در زنجیره توزیع و چه در زنجیره تأمین میزان کالا و حتی کیفیت و نوع کالا را هم پیش‌بینی کند. به طور یقین در ادامه اجرای این

مدل تغییر رویکرد و انتظارات پنهان و کشف نشده مشتریان نهایی را هم قابل پیش‌بینی خواهد نمود.

این مقاله مبتنی بر یکی از روش‌های یادگیری عمیق در هوش مصنوعی به نام شبکه حافظه طولانی کوتاه‌مدت^۱ LSTM صورت می‌پذیرد. برای پیش‌بینی موجودی انبار توسط کلان داده‌های صنعتی، با رویکرد راهبردی، ابتدا بر طراحی الگویی در زمینه شناسایی ابعاد، عناصر، مؤلفه‌ها و شاخص‌های اصلی جهت همگرایی تحلیل کلان داده و پیش‌بینی نقطه بهینه سفارش انبار متمرکز خواهد گردید. رهیافت و پارادایم این مقاله ضمن بررسی رویکرد اثبات گرایانه که با نگاه از بیرون در قالب اعداد و ارقام، تئوری‌ها و معروف‌ترین مدل‌های هوش مصنوعی، با رویکرد تفسیرگرایی، از طریق تحلیل روابط داده‌ای به رویداد گرایی مبتنی بر معنا بخشی به واقعیات و همچنین در تعامل با موضوع مقاله، به ایجاد و توسعه فرضیات و مدل‌های تئوریک ناشی از ساخت‌های پنهان و فرایندها خواهد پرداخت. از این‌روی، با توجه به مطالب فوق، این مقاله با به‌کارگیری یک رویکرد ابتکاری هوش مصنوعی در زمینه کنترل موجودی انجام شده است. در پایان مدل رویکرد ترکیبی پیشنهادی بر روی داده‌های محیط واقعی پیاده‌سازی شده است.

۲- پیشینه تحقیق

همان‌گونه که توسط مایر و همکاران ([Mayer & Schoenberger, 2013](#))، در مقاله‌ای عنوان شد، تعریف رسمی و دقیق از کلان داده وجود ندارد ([Mayer-Schoenberger & Cukier, 2013](#) & شی [Shi, 2014](#))، کلان داده را به دلیل پیچیدگی، تنوع و نامتجانس بودن برای پردازش و تحلیل در زمانی معقول، دشوار معرفی کرد درعین‌حال که معتقد بود کلان داده از نگاه سیاست‌گذاران بسیار باارزش و منبعی راهبردی و کلیدی جهت رهبری تحول در عصر فناوری و پیش رو بودن در نوآوری می‌تواند باشد ([Shi, 2014](#)) بعدها صنایعی ([Sanaei, 2016](#)) در مقاله خود بر این عقیده بود که کلان داده عامل اکتشافات اصول و روابطی است که تنها با وجود مجموعه داده‌های بزرگ میسر خواهند بود. در ادامه مارو و همکاران ([Mauro et al, 2016](#)) چنین یافتند که کلان داده شامل اطلاعاتی است که

¹ Long short-Term Memory

به واسطه تنوع داده‌ها، حجم و سرعت تغییر آنها، برای ارزشمند شدن نیاز به ابزارهای فناورانه و روش‌های تحلیلی خاص دارد.

در مرحله برنامه‌ریزی مدیریت انبارهای گسترده، کلان داده، نقش حیاتی دارد زیرا کمک زیادی به شرکت‌ها برای تصمیم‌گیری استراتژیک در زمینه منبع‌یابی، طراحی شبکه زنجیره تأمین و همچنین طراحی و توسعه محصول می‌کند. همچنین در مرحله عملیات زنجیره تأمین، کلان داده، به تصمیم‌گیری در عملیات زنجیره تأمین کمک می‌کند مثل برنامه‌ریزی تقاضا، تدارکات، تولید، موجودی انبار و حمل‌ونقل (Wamba et al, 2016). اصطلاح «کلان داده‌ها» به مجموعه داده‌هایی اطلاق می‌شود که از نظر سرعت، حجم و تنوع در سطح بالایی هستند و با تکنیک‌ها و ابزارهای سنتی پردازش نمی‌شوند. در حال حاضر، کلان داده‌ها همه‌جا هستند، چه به شکل داده‌های ساخت‌یافته، مانند پایگاه داده‌های سنتی سازمان (مثلاً سیستم مدیریت ارتباط با مشتری) یا داده‌های بدون ساختار که فناوری‌های جدید ارتباطی و بسترهایی هستند که کاربر می‌تواند آنها را توسعه دهد یا ویرایش کند؛ مثلاً متون، تصاویر و فیلم‌ها (Motaharinejad et al, 2017).

در ساده‌ترین حالت، یادگیری عمیق را می‌توان راهی برای خودکارسازی تجزیه و تحلیل پیشگویانه در نظر گرفت (Wazan, 2022). یادگیری عمیق یک زیرشاخه خاص از یادگیری ماشین است، برداشتی جدید از بازنمایی‌های یادگیری از داده‌ها که بر یادگیری لایه‌های متوالی به طور فزاینده تأکید دارد (Shao, 2020). اکثر الگوریتم‌های یادگیری ماشین حاوی فرآیندهایی هستند که باید بیرون از خود الگوریتم یادگیری و با استفاده از داده‌های اضافی تعیین گردند. یادگیری ماشین در واقع نوعی از آمار کاربردی است که در آن، تمرکز بیشتری بر استفاده از رایانه برای تخمین توابع پیچیده و تمرکز کمتری بر روی اثبات فواصل اطمینان حول این توابع صورت می‌گیرد (Tarabian, 2017). اسمیت و همکاران (Smith et al, 2021) بر این عقیده هستند که یادگیری عمیق شامل مجموعه‌ای از مدل‌های یادگیری ماشینی است که در دو حالت یادگیری با ناظر و بدون ناظر در معماری سلسله مراتبی عمیق قرار می‌گیرند. در واقع، در شبکه‌های عصبی یادگیری با ناظر (طبقه‌بندی) ورودی، خروجی را طبقه‌بندی کرده و بر اساس داده‌های قبلی که برچسب دار هستند می‌توان داده‌های جدید را پیش‌بینی کرد (Smith et al, 2021). هدف یادگیری عمیق این است که ویژگی‌های یادگیری سلسله مراتبی سطح بالا را از ویژگی‌های

سطح پایین یاد بگیرد، یعنی در لایه‌های ابتدایی ویژگی‌های ساده مثلاً لبه‌ها و خط‌ها و در ویژگی‌های میانی گوشه‌ها، لبه‌ها و سپس به ترتیب ویژگی‌های سطح بالاتر یاد گرفته شود (Goodfellow et al, 2013) اکثر روش‌های یادگیری عمیق از معماری شبکه عصبی استفاده می‌کنند، به همین دلیل اغلب از مدل‌های یادگیری عمیق به‌عنوان شبکه عصبی عمیق یاد می‌شود. اصطلاح «عمیق» معمولاً به تعداد لایه‌های پنهان در شبکه عصبی اشاره دارد. شبکه‌های عصبی سنتی تنها شامل ۲ تا ۳ لایه پنهان هستند، درحالی‌که شبکه‌های عصبی عمیق می‌توانند تا ۱۵۰ لایه داشته باشند. مدل‌های یادگیری عمیق با استفاده از مجموعه‌های بزرگی از داده‌های برجسب‌گذاری شده و معماری شبکه عصبی، آموزش داده می‌شوند که ویژگی‌ها را مستقیماً از داده‌ها، بدون نیاز به استخراج ویژگی به‌صورت دستی، یاد می‌گیرند (Goodfellow et al, 2016) در مقاله‌ای بر این باورند که ارائه و حل یک مدل ریاضی به‌منظور کنترل موجودی متناسب با فضای صنعت، در راستای کمینه‌کردن هزینه‌های موجودی و به‌دست‌آوردن مقدار بهینه سفارش ضروری است و توجه به فضای قفسه برای کنترل موجودی کالاهای فسادپذیر، بدون در نظر گرفتن فضای انبار امر مرسوم است (Riyazi et al, 2022). آنها با این هدف که در کنار توجه به فروش کالای پُرسودتر و عقد قرارداد حداقل تعهد خرید کالا با تأمین‌کنندگان، به‌نحوی که هزینه کل دوره سفارش‌گذاری بهینه شود، به انجام این مقاله پرداختند و مدلی برای کنترل موجودی با محدودیت‌های جدید ارائه کردند. همچنین در فرایند حل مسئله، برای به‌دست‌آوردن جواب بهینه، از الگوریتم بهینه‌سازی ازدحام ذرات به کمک نرم‌افزار متلب و برای پیش‌بینی تقاضا، از سری زمانی در نرم‌افزار مینی‌تب استفاده کرده در مقاله‌ای نیز بر این باورند، پذیرش کانال‌های آنلاین و تجارت الکترونیک، به تغییرات مداوم و پویا در صنعت خرده‌فروشی منجر شده که یک توسعه اجتناب‌ناپذیر است و بسیاری از شرکت‌ها را با چالش انتخاب مناسب‌ترین کانال فروش، برای ارائه یک تجربه یکپارچه به مشتریان خود مواجه کرده است. خرده‌فروشی همه‌جانبه یکپارچه، با مفهوم ادغام همه کانال‌ها، ضمن ایجاد تجربه مذکور، باعث افزایش پیچیدگی فرآیندهای پیش‌بینی و برنامه‌ریزی می‌شود (Sultani et al, 2023).

فیاضی و همکاران (Faizi Rad et al, 2021) در مقاله‌ای برای پیش‌بینی تقاضا در سیستم‌های رزرواسیون دانشگاهی با هدف کاهش ضایعات مواد غذایی به کمک شبکه‌های

عصبی، جهت پیش‌بینی تقاضای واقعی، از شبکه عصبی مصنوعی با تابع خطای موزونی که به کمک جست‌وجوی الگوی تعمیم‌یافته جهت‌دهی می‌شود، استفاده شد. شاخص‌های مجموع رزرو، روز هفته، سطح قیمت وعده، مجموع تعداد رزرو، تعداد رزرو به تفکیک مقطع تحصیلی، تعداد رزرو به تفکیک وضعیت اسکان و غذای مجاور به‌عنوان متغیرهای ورودی و تعداد تقاضای واقعی غذا نیز شاخص خروجی در نظر گرفته شد. آنها داده‌های هفت سال اخیر سامانه رزرواسیون سلف مرکزی یکی از دانشگاه‌های بزرگ کشور که سالانه به‌طور متوسط پتانسیل تولید ۵۶ هزار پرس غذای مازاد (بیش از ۲۳ هزار تن مواد غذایی) را دارد، بررسی کردند. با آموزش یک شبکه عصبی مصنوعی و الگوریتم ترکیبی با تابع خطای موزون متناسبی را به‌دست آوردند که قادر است تولید روزانه غذای مازاد را بیش از ۸۰ درصد کاهش دهد.

توجه همزمان به زمان و میزان سفارش کالا و به حداقل رساندن سیستم و هزینه‌های مشتری از دغدغه‌های اصلی مدیریت موجودی است. پونیا و همکاران ([Punia et al., 2020](#)) در مقاله خود با عنوان «یک چارچوب سلسله مراتبی زمانی و یادگیری عمیق برای پیش‌بینی زنجیره تأمین»، یک چارچوب پیش‌بینی زمانی جدید که به شکل سلسله مراتبی است، برای تولید پیش‌بینی‌های منسجم تمام سطوح زنجیره تأمین خرده‌فروشی پیشنهاد کرده‌اند. آنها در ابتدا با استفاده از روش پیش‌بینی پایه (شبکه‌های LSTM) پیش‌بینی‌های مستقیم (مجموع ۱۱۰ سری) انجام داده و برای همین تعداد از سری‌های زمانی، پیش‌بینی با استفاده از چارچوب سلسله مراتبی زمانی را محاسبه کردند. سپس $ARMAE$ و $ARMSE$ را برای پیش‌بینی‌های کوتاه‌مدت، میان‌مدت و بلندمدت محاسبه و به این نتیجه رسیدند که پیش‌بینی‌های چارچوب پیشنهادی به‌طور قابل‌توجهی بهتر از پیش‌بینی‌های مستقیم هستند. علاوه بر این، بهبودها در سطوح مقطعی و زمانی زنجیره تأمین معنادار و ثابت هستند. علاوه بر این، مشاهده شد که پیش‌بینی‌های پایین به بالا نسبت به پیش‌بینی‌های بالا به پایین، زمانی که از داده‌های نقطه‌ای فروش برای پیش‌بینی در زنجیره تأمین خرده‌فروشی آنلاین و آفلاین استفاده می‌شود، دقیق‌تر هستند.

نیکولوپولوس و همکاران ([Nikolopoulos et al., 2020](#)) در مقاله خود با عنوان «پیش‌بینی و برنامه‌ریزی در طول یک بیماری همه‌گیر: نرخ رشد کوید ۱۹، اختلالات زنجیره تأمین و تصمیمات دولتی» با استفاده از داده‌های ایالات متحده، هند، بریتانیا، آلمان و سنگاپور تا

اواسط آوریل ۲۰۲۰، نرخ رشد کوید ۱۹ را با مدل‌های آماری، اپیدمیولوژی و یادگیری عمیق پیش‌بینی و یک روش جدید پیش‌بینی ترکیبی بر اساس نزدیک‌ترین همسایه‌ها و خوشه‌بندی پیشنهاد کردند. در ادامه با استفاده از داده‌های کمکی (روند گوگل) و شبیه‌سازی تصمیمات دولتی، تقاضای اضافی محصولات و خدمات در طی این بیماری همه‌گیر پیش‌بینی می‌شود. نتایج تجربی آنها به سیاست‌گذاران و برنامه‌ریزان کمک کند تا تصمیمات بهتری را طی دوران بیماری همه‌گیر اتخاذ کنند.

فان و همکاران (Fan et al, 2021)، در مقاله خود با عنوان «تحقیق در مورد چارچوب موجودی سریع بر اساس شبکه‌های عصبی عمیق»، فناوری هوش مصنوعی برای موجودی مربوط به قفسه‌های کتاب بر اساس شبکه عصبی عمیق پیشنهاد دادند. این چارچوب از برجسب‌های بارکد و عکس‌ها یا جریان‌های قفسه برای درک موجودی و مرتب‌سازی مؤثر کتاب‌ها استفاده می‌کند. در مقایسه با سیستم موجودی‌ای که از فناوری RFID استفاده کردند. حاصل یافته‌های آنها این بود که این فناوری دارای مزیت‌هایی همچون هزینه پیاده‌سازی پایین و سرعت بالا است. فناوری موجودی مبتنی بر شبکه عصبی عمیق تنها نیازمند نصب برنامه موجودی بر روی یک سرور است. سپس، کتابداران از قفسه کتاب با تلفن یا دوربین عکس می‌گیرند و عکس‌ها را به برنامه آپلود می‌کنند. همچنین، فناوری آماربرداری مبتنی بر شبکه عصبی عمیق می‌تواند عکس‌های قفسه را برای دوره‌های طولانی حفظ کند و در صورت نبود کتاب، کتابداران می‌توانند با بازیابی تصاویر، کتاب‌های گم‌شده را بررسی کنند.

فیچر و کراس (Fischer & Krauss, 2018)، در مقاله خود با عنوان «یادگیری عمیق با شبکه‌های حافظه بلندمدت برای پیش‌بینی‌های بازار مالی» با استفاده از شبکه‌های LSTM، روش جنگل تصادفی، یک شبکه عصبی عمیق (DNN) و یک رگرسیون لجستیک، الگوهای مشترکی را در سهام برتر و شکست‌خورده استخراج می‌کنند و منابع سودآوری را آشکار می‌سازند. نهایتاً یک استراتژی ساده تجاری را بر اساس این یافته‌ها توسعه می‌دهد. این مقاله از دسامبر ۱۹۹۲ تا اکتبر ۲۰۱۵، شبکه‌های حافظه بلندمدت را به یک وظیفه پیش‌بینی بازار مالی در مقیاس بزرگ اعمال می‌کند و روی کاربرد تجربی بزرگ مقیاس شبکه‌های LSTM در وظایف پیش‌بینی سری‌های زمانی مالی تمرکز دارد.

کومار و همکاران (Kumar et al, 2020)، در مقاله خود برای پیش‌بینی تقاضامحور با تأثیر متغیرهای آمیخته بازاریابی و با هدف بهبود دقت پیش‌بینی تقاضا، یک مدل شبکه عصبی مبتنی بر انتشار معکوس با ورودی‌های فازی آموزش می‌دهند و با روش‌های پیش‌بینی معیار بر روی داده‌های سری زمانی با استفاده از داده‌های تقاضای تاریخی و فروش در ترکیب با اثربخشی تبلیغات، هزینه و تبلیغات مقایسه می‌کنند. تحلیل آماری انجام‌شده و آزمایش‌ها نشان می‌دهد که روش استفاده‌شده در چارچوب پیشنهادی در بهینگی، کارایی و سایر معیارهای آماری بهتر عمل می‌کند. در نهایت، برای بهبود دقت پیش‌بینی شبکه‌های عصبی فازی، توسعه برنامه‌های بازاریابی برای محصولات و بحث در مورد پیامدهای آن‌ها در زمینه‌های مختلف، دیدگاه‌های ارزشمندی برای مدیران ارائه می‌کند.

کیم و همکاران (Kim et al, 2022)، در مقاله خود با عنوان «چارچوب پیش‌بینی مبتنی بر 2D KDE و LSTM برای مدیریت موجودی مقرون‌به‌صرفه در تولید هوشمند» بر این باورند، در اغلب شرکت‌های کوچک و متوسط، مدیریت سیستماتیک داده‌ها وجود ندارد و به دلیل فقدان داده‌ها و نوسانات داده‌های تصادفی، پیش‌بینی مدل‌ها به‌خوبی کار می‌کند. از آنجاکه مدل پیش‌بین یک تابع اصلی مشتق شده از مدیریت داده‌های موجودی شرکت است، عملکرد ضعیف مدل موجب تخریب سیستم مدیریت داده‌های موجودی شرکت می‌شود. آنها در مقاله خود چارچوبی برای پیش‌بینی قابل‌اطمینان داده‌های موجودی یک شرکت با مدل‌سازی نوسانات یک شرکت به‌صورت تصادفی ارائه کرده‌اند و چارچوب پیش‌بینی را با استفاده از مدل پیش‌بینی نقطه‌ای با استفاده از LSTM (حافظه کوتاه‌مدت بلندمدت)، تابع چگالی کرنل سه‌بعدی و نتایج پیش‌بینی هزینه مدیریت موجودی بررسی کردند. با انجام آزمایش‌های مختلف، ضرورت پیش‌بینی فاصله در پیش‌بینی تقاضا و اعتبار مدل پیش‌بینی مؤثر هزینه از طریق تابع readjustment نشان داده شده است. حاصل یافته‌های آنان، قابلیت اطمینان نتایج پیش‌بینی را تضمین می‌کند و الگوریتم توسعه یافته در این مقاله، اثر صرفه‌جویی هزینه را با احتمال ۹۸٫۸۷ درصد در آزمایش داده‌های واقعی نشان داد.

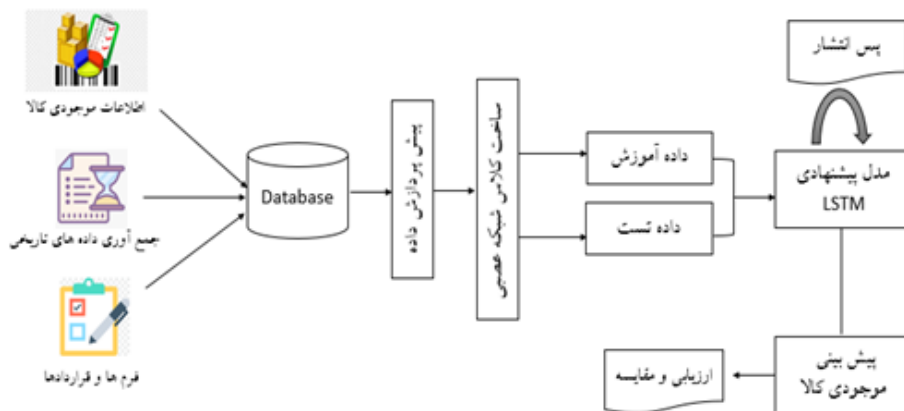
جهت تحقق اهداف این مقاله از شبکه عصبی حافظه‌ی کوتاه‌مدت طولانی (LSTM) که نوعی خاص از شبکه عصبی بازگشتی (RNN) محسوب می‌شود استفاده شده است. علت استفاده از LSTM به‌منظور جلوگیری از مشکل وابستگی طولانی‌مدت در شبکه است. رویکرد مقاله، از نوع روش‌شناسی / رهیافت کمی کیفی یا ترکیبی خواهد بود، لذا لازم است هر دو نوع داده کمی و کیفی جمع‌آوری شود. تحلیل داده‌های کیفی برگرفته از پیشینه کاوی، پس از کدگذاری بر اساس علم آمار و تحلیل داده‌های کمی با استفاده از روش یادگیری عمیق از حوزه یادگیری ماشین، برای پاسخ به سؤالات، مورد استفاده قرار می‌گیرد. استفاده از روند هم‌زمانی، جهت همگرایی اطلاعات کمی و کیفی جمع‌آوری شده، مد نظر است. این مقاله با استفاده از زبان برنامه‌نویسی پایتون در محیط کولب نوشته و اجرا می‌شود. جهت پیاده‌سازی شبکه‌های عصبی و مدل‌سازی داده‌ها از کتابخانه pandas جهت کار با مجموعه داده‌ها و عملیات مختلف، کتابخانه numpy جهت کار با آرایه‌ها و محاسبات مختلف و عملیات ماتریسی، کتابخانه seaborn جهت مصورسازی داده‌ها، کتابخانه matplotlib جهت نمایش داده‌ها به‌صورت گرافیکی و همچنین از کتابخانه scikit learn که یکی از معروف‌ترین و پرکاربردترین کتابخانه‌های داده‌کاوی و یادگیری ماشین در پایتون است، جهت استفاده از الگوریتم‌های پیش‌پردازش، طبقه‌بندی، خوشه‌بندی و کاهش ابعاد مورد استفاده قرار می‌گیرد. نهایتاً از imblearn برای کار با داده‌های نامتوازن و کتابخانه OS جهت تعامل با سیستم استفاده خواهیم کرد. همچنین جهت ارزیابی مجموعه داده‌ها از نمودارهای متغیر عددی با تخمین چگالی هسته و هیستوگرام استفاده خواهد شد و نتایج با استفاده از ماتریس اغتشاش ارائه خواهد گردید.

این مقاله طی مراحل زیر انجام شده است:

مرحله اول: بررسی و مطالعه فرآیند زنجیره تأمین و چالشهای موجود برای تهیه محصولات مورد نیاز بازار با توجه به شرایط و میزان فروش.

مرحله دوم: داده‌های آموزشی برای پیش‌بینی محبوبیت و میزان فروش محصول بسیار مهم هستند. از این رو، ما باید اطلاعات با کیفیت بالا را جمع‌آوری کنیم. همچنین می‌توانیم به‌طور منظم اطلاعات سفارش محصول را از فروشگاه‌ها جمع‌آوری کنیم. این اطلاعات می‌توانند مستقیماً برای منعکس کردن توزیع نیازهای مشتری به میزان معینی مورد استفاده قرار گیرند.

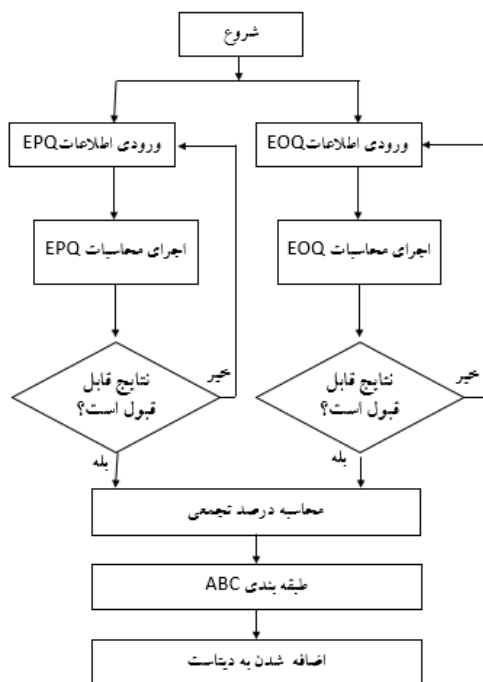
جهت انجام این مقاله و ارزیابی رویکرد پیشنهادی از مجموعه داده های موجود در سایت کگل (kaggle) که در خصوص یک پایگاه داده به تعداد ۳۹۷۸۴ رکورد بوده استفاده شده که مربوط به کنترل موجودی یک خرده فروشی است و برای استفاده در پروژه های تحقیقاتی در اختیار عموم پژوهشگران قرار داده شده ، استفاده گردیده است. لازم به ذکر است این مجموعه داده در تاریخ ۱۳ اکتبر ۲۰۲۳ بروزرسانی شده است. به طور کلی، از این مجموعه داده می توان برای تجزیه و تحلیل و مدیریت موجودی محصولات مختلف استفاده کرد. اطلاعاتی در مورد تاریخچه فروش، میزان سفارش، هزینه نگهداری، هزینه فرصت، هزینه سفارش، قیمت کالا، نوع بازاریابی، جزئیات انتشار، فواصل تکرار نوع بازار، فواصل تکرار نوع محصول، فواصل تکرار شونده محصول جدید، فواصل تکرار نوع محصول و وضعیت محصولات فعال ارائه می دهد. برای اجرای مدل ۸۰ درصد دادگان جهت آموزش و ۲۰ درصد برای تست مدل در نظر گرفته شده است. رویکرد اتخاذ شده به شدت به ماهیت، کمیت و کیفیت داده ها موجود وابسته است، بنابراین ارزیابی داده ها و انتخاب داده هایی که دارای کیفیت مطلوب باشند، یک مرحله مهم و اساسی به حساب می آید. شکل زیر روند کلی طرح پیشنهادی و گام های اجرا را نمایش می دهد.



شکل ۱: جریان کلی رویکرد پیشنهادی

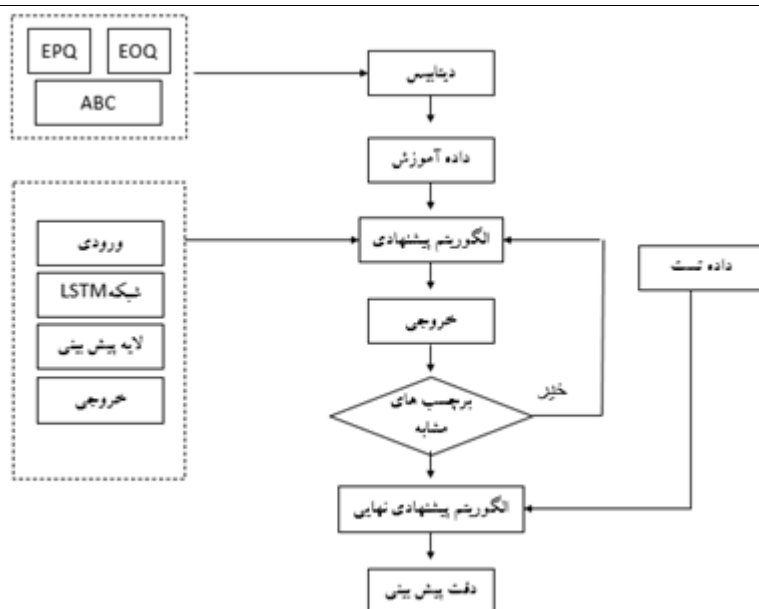
مرحله سوم: پیش پردازش داده ها و نرمال سازی داده ها جهت مدلسازی رویکرد پیشنهادی گامی مهم به شمار می آید از این رو، در این مرحله اختلالات پیش پردازش داده ها در راستای حذف داده های پرت و غیر نرمال، حذف و انتخاب ویژگیهای موثر میباشد که قبل از پیاده سازی مدل باید انجام شود

مرحله چهارم: رویکردهای سنتی کنترل موجودی با تکنیک‌های یادگیری عمیق ترکیب می‌شوند تا داده‌های جدید جهت پیش‌بینی به شبکه پیشنهادی تزریق شوند. در راستای اهداف این پژوهش، ابتدا از دو تکنیک کنترل موجودی EOQ و EPQ برای ایجاد یک طبقه‌بندی باینری استفاده می‌کنیم تا لیستی از شناسه محصول را داشته باشیم و موجودی یا لیست محصولاتی که باید حذف شوند، نگهداری شوند. سپس با محاسبه درصد تجمعی محصولات فروخته شده، تجزیه و تحلیل ABC با استفاده از اصل پارتو صورت می‌گیرد.



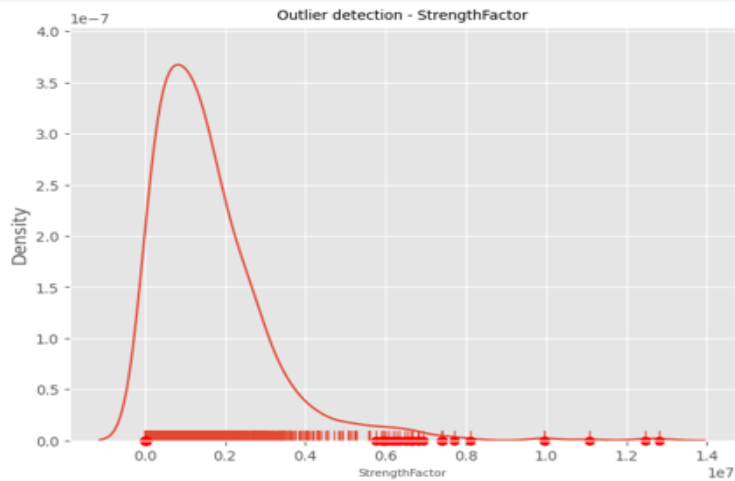
شکل ۲: دیاگرام مربوط به روند دیتاسازی

مرحله پنجم: پیاده سازی مدل و اجرای رویکرد پیشنهاد شده. از آنجاکه پیدایش هوش مصنوعی و تکنیک‌های یادگیری عمیق، به عنوان رویکردی نوین، پتانسیل بالایی در کاربردهای مختلف ارزیابی و تحلیل از خود نشان داده است، جهت پیش‌بینی موجودی کالا از شبکه عصبی LSTM استفاده می‌کنیم. دلیل انتخاب این دو شبکه قابلیت یادگیری وابستگی‌های طولانی مدت است. مراحل انجام کار با استفاده از دیتاستی که هم شامل داده‌های تاریخی فروش و هم موجودی حاضر هستند، انجام شده است.



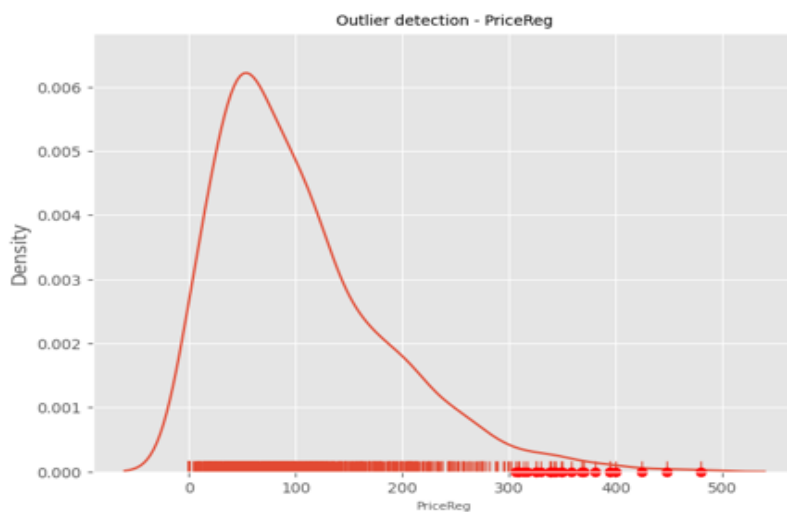
شکل ۳: دیاگرام روند پیاده‌سازی رویکرد پیشنهادی

مرحله ششم: عملکرد مدل پیشنهاد شده با روشهای ارزیابی بررسی میشود. دقت و کارایی رویکرد پیشبینی محاسبه می گردد، و معمولاً مقایسات دارای نتایج معناداری خواهند بود. جهت پیش‌پردازش داده‌ها، نقاط پرت تک متغیره را برای تجزیه و تحلیل نقاط پرت در ویژگی‌های عددی مجموعه داده مشخص می‌کنیم. نقاط پرت در داده‌های ورودی می‌تواند ما را همراه کند و نتایج قابل اعتمادی حاصل نکند؛ بنابراین برای رفع داده‌های پرت تک متغیره بسیاری از الگوریتم‌ها به محدوده و توزیع مقادیر ویژگی در داده‌های ورودی حساس هستند. نقاط و داده‌های پرت در داده‌های ورودی می‌توانند نتایج را منحرف کنند و نتایج را کمتر قابل اعتماد کنند، به همین دلیل است که ما باید همه موارد پرت را بشناسیم و آنها را رفع کنیم.



شکل ۴: نمودار تشخیص پراکندگی قدرت عامل‌ها

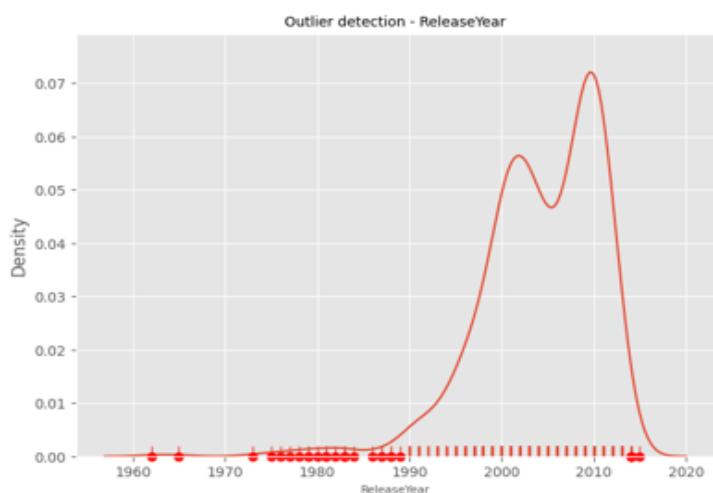
محور افقی (X): نشان‌دهنده ضریب قدرت محصولات. این ویژگی به احتمال تشخیص آنها به عنوان نقاط پرت مربوط است. محور عمودی (Y): دانسیته یا تراکم نقاط پرت شناسایی شده با توجه به ضریب قدرت. این نمودار نشان می‌دهد که بیشتر نقاط پرت در محدوده ضریب قدرت ۰,۵ تا ۱,۰ قرار دارند.



شکل ۵: نمودار تشخیص پراکندگی قیمت محصولات

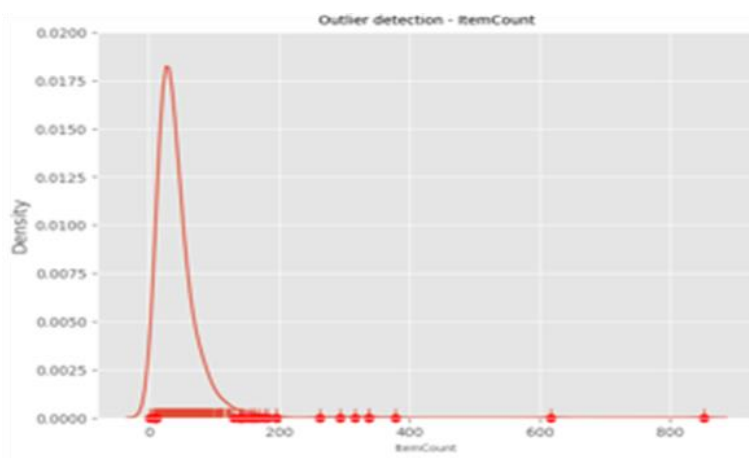
محور افقی (X): نشان‌دهنده قیمت محصولات. این ویژگی به احتمال تشخیص آنها به عنوان نقاط پرت مربوط است. محور عمودی (Y): دانسیته یا تراکم نقاط پرت شناسایی شده با

توجه به قیمت محصولات. این نمودار نشان می‌دهد که بیشتر نقاط پرت در محدوده قیمت ۰ تا ۴۰۰ قرار دارند.



شکل ۶: نمودار تشخیص پراکندگی سال انتشار

محور افقی (X): نشان‌دهنده سال انتشار محصولات. این ویژگی به احتمال تشخیص آنها به عنوان نقاط پرت مربوط است. محور عمودی (Y): دانسیته یا تراکم نقاط پرت شناسایی شده با توجه به سال انتشار محصولات. این نمودار نشان می‌دهد که بیشتر نقاط پرت در محدوده سال‌های ۱۹۶۰ تا ۱۹۸۰ قرار دارند.



شکل ۷: نمودار تشخیص پراکندگی تعداد آیتم‌ها

محور افقی (X): نشان‌دهنده تعداد آیتم‌ها در هر محصول. این ویژگی برای تشخیص نقاط پرت در نظر گرفته شده است. محور عمودی (Y): دانسیته یا تراکم نقاط پرت شناسایی شده با توجه به تعداد آیتم‌ها. این نمودار نشان می‌دهد که بیشتر نقاط پرت در محدوده ۲۰۰ تا ۶۰۰ آیتم قرار دارند.



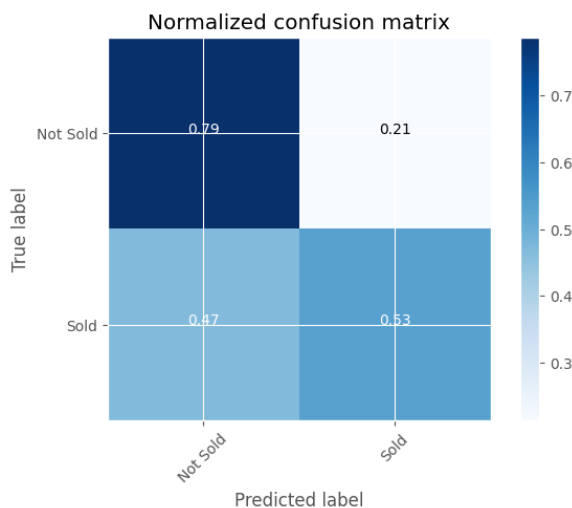
شکل ۸: نمودار تشخیص پراکندگی قیمت

محور افقی (X): نشان‌دهنده قیمت پایین کاربران (LowUserPrice) برای هر محصول. این ویژگی برای تشخیص نقاط پرت در نظر گرفته شده است. محور عمودی (Y): دانسیته یا تراکم نقاط پرت شناسایی شده با توجه به قیمت پایین کاربران. این نمودار نشان می‌دهد که بیشتر نقاط پرت در محدوده قیمت ۰ تا ۰,۰۰۸ قرار دارند.

پیش‌بینی مازول طبقه‌بندی و نرمال‌سازی داده‌ها

مدل پیش‌بینی کننده مازول طبقه‌بندی شناسه منحصر به فرد هر محصول را پیش‌بینی می‌کند که باید در موجودی نگهداری شود؛ یعنی همان وضعیت فعال. سپس در مرحله بعدی با استفاده از تابع SMOTH کلاس‌ها متعادل می‌گردند. نهایتاً با استفاده از ماتریس سردرگمی نرمال شده و نشده را محاسبه و ترسیم می‌گردد. برای ارزیابی عملکرد مدل پیش‌بینی موجودی بسازید. این ماتریس معمولاً در مسائل دسته‌بندی استفاده نموده‌ایم و نشان می‌دهیم که مدل چه میزان داده‌ها را به درستی و به اشتباه دسته‌بندی کرده است. ابتدا تابع `plot_confusion_matrix` تعریف شده است. این تابع یک ماتریس سردرگمی، نام

کلاس‌ها و دیگر پارامترهای مرتبط را دریافت می‌کند و نمودار ماتریس سردرگمی را رسم می‌کند. سپس، ماتریس گیجش محاسبه می‌شود با استفاده از تابع `confusion_matrix` که از کتابخانه `sklearn.metrics` استخراج می‌شود. این تابع دو آرگومان می‌پذیرد `testing_target` که برچسب‌های واقعی یا حقیقی برای داده‌ها است و `pred` که پیش‌بینی‌های مدل برای همان داده‌ها است. ماتریس سردرگمی با استفاده از این دو مقدار محاسبه می‌شود. سپس، دو نمودار ماتریس سردرگمی رسم می‌شوند. اولین نمودار، بدون نرمال‌سازی است و دومین نمودار با نرمال‌سازی است. نرمال‌سازی ماتریس سردرگمی می‌تواند مفید باشد زیرا نشان می‌دهد که هر کلاس چه میزان درست تشخیص داده شده است نسبت به کل تعداد نمونه‌های آن کلاس در کل دیتاست را نمایش می‌دهد. در نمودارها، رنگ نقاط مشخص می‌کند که هر دسته چه میزان درست یا نادرست تشخیص داده شده است. همچنین، اعداد درون هر نقطه نشان‌دهنده مقدار درصد تعداد نمونه‌های هر کلاس است که به درستی یا نادرست تشخیص داده شده‌اند. استفاده از ماتریس سردرگمی و نمودارهای مربوطه می‌تواند به ما کمک می‌کند تا عملکرد مدل خود را در پیش‌بینی موجودی بسنجیم و در صورت نیاز اقدامات بهبود را انجام دهیم.



شکل ۹: ماتریس سردرگمی نرمال شده

ماتریس تعمیم‌یافته آشفته‌گی ۱ که در تصویر نشان داده شده است، عملکرد یک مدل طبقه‌بندی را ارزیابی می‌کند. این ماتریس برای نشان دادن تعداد نمونه‌های طبقه‌بندی شده صحیح و نادرست در هر کلاس استفاده می‌شود.

مفاهیم کلیدی دقیق و مورد نیاز:

کلاس: هر ردیف و ستون در ماتریس، یک کلاس را نشان می‌دهد. در این تصویر، دو کلاس وجود دارد: «فروخته نشده» و «فروخته شده».

نمونه: هر سلول در ماتریس، تعداد نمونه‌ها را در یک دسته خاص نشان می‌دهد. به عنوان مثال، سلول (۱, ۱) تعداد نمونه‌هایی را نشان می‌دهد که به درستی به عنوان «فروخته نشده» طبقه‌بندی شده‌اند.

دقت: دقت کلی مدل، نسبت تعداد نمونه‌های طبقه‌بندی شده صحیح به کل نمونه‌ها است. در این تصویر، دقت ۰,۶۷ است، به این معنی که ۶۷٪ از نمونه‌ها به درستی طبقه‌بندی شده‌اند.

حساسیت: حساسیت برای هر کلاس، نسبت تعداد نمونه‌های واقعی مثبت که به درستی به عنوان مثبت طبقه‌بندی شده‌اند به کل نمونه‌های مثبت واقعی است. در این تصویر، حساسیت برای کلاس «فروخته شده» ۰,۴۷ و برای کلاس «فروخته نشده» ۰,۷۹ است.

اختصاصیت: اختصاصیت برای هر کلاس، نسبت تعداد نمونه‌های واقعی منفی که به درستی به عنوان منفی طبقه‌بندی شده‌اند به کل نمونه‌های منفی واقعی است. در این تصویر، اختصاصیت برای کلاس «فروخته شده» ۰,۵۳ و برای کلاس «فروخته نشده» ۰,۷۱ است. در این ماتریس، ۷۹٪ از نمونه‌هایی که واقعاً فروخته نشده بودند، به درستی به عنوان «فروخته نشده» طبقه‌بندی شده‌اند.

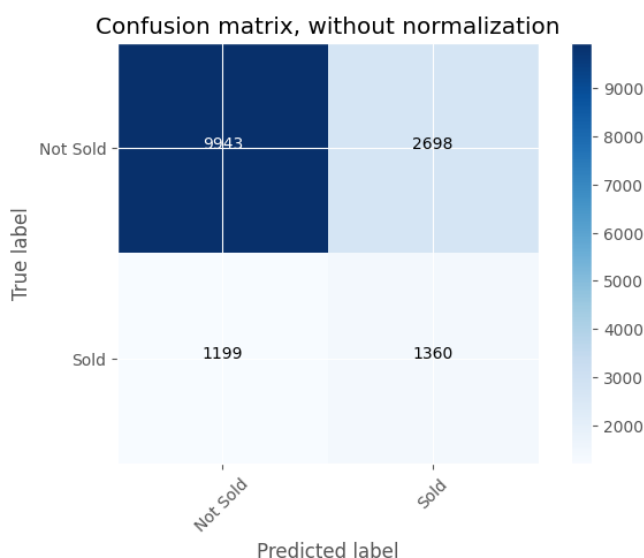
۲۱٪ از نمونه‌هایی که واقعاً فروخته نشده بودند، به اشتباه به عنوان «فروخته شده» طبقه‌بندی شده‌اند.

۴۷٪ از نمونه‌هایی که واقعاً فروخته شده بودند، به درستی به عنوان «فروخته شده» طبقه‌بندی شده‌اند.

۵۳٪ از نمونه‌هایی که واقعاً فروخته شده بودند، به اشتباه به عنوان «فروخته نشده» طبقه‌بندی شده‌اند.

¹ Normalized confusion matrix

این ماتریس تعمیم‌یافته آشفتگی نشان می‌دهد که مدل طبقه‌بندی در تشخیص نمونه‌های فروخته نشده عملکرد بهتری نسبت به تشخیص نمونه‌های فروخته شده دارد. با این حال، هنوز جای پیشرفت وجود دارد، به‌خصوص در مورد کاهش تعداد نمونه‌هایی که به اشتباه طبقه‌بندی شده‌اند. ماتریس تعمیم‌یافته آشفتگی فقط یک ابزار برای ارزیابی عملکرد مدل است.



شکل ۱۰: ماتریس سردرگمی نرمال نشده

این ماتریس نشان‌دهنده عملکرد یک مدل طبقه‌بندی است. اعداد در سلول‌ها نشان‌دهنده تعداد نمونه‌هایی است که در هر دسته قرار گرفته‌اند. سطرها نشان‌دهنده برچسب‌های واقعی (true label) و ستون‌ها نشان‌دهنده برچسب‌های پیش‌بینی شده توسط مدل predicted label است. عدد ۹۸۴۳ در سلول (Not Sold, Not Sold) یعنی ۹۸۴۳ نمونه به درستی به عنوان Not Sold پیش‌بینی شده‌اند. عدد ۱۱۹۹ در سلول (Sold, Sold) یعنی ۱۱۹۹ نمونه به درستی به عنوان Sold پیش‌بینی شده‌اند.

دیتاسازی جهت تزریق به شبکه

در این مرحله قصد داریم تعیین کنیم که کدام محصولات از موجودی خود را برای فروش نگه داریم و کدام محصولات را که فروش خوبی ندارند را کنار بگذاریم. لازم به ذکر است داده‌ها هم شامل داده‌های تاریخی فروش (داده‌های باسابقه یا غیرفعال) و هم موجودی فعال هستند. به ازای هر سطر از دادگان محاسبه مقدار سفارش اقتصادی (EOQ) و مدل تولید اقتصادی (EPQ) و همچنین محاسبه مقادیر آنالیز ABC نیز افزون بر داده‌های اولیه نیز با استفاده از زبان برنامه‌نویسی پایتون بر روی کل دادگان اعمال شده است و در نهایت ۳ ستون جدید به دادگان اضافه شده است. بدیهی است جدول دادگان بروز رسانی شده به‌صورت ذیل تغییر یافته است. لازم به ذکر است، این فرآیند طبق عملکرد کد پایتون زیر صورت پذیرفته است و محاسبات جدید منجر به بروز رسانی جدول شده است.

EOQ	EPQ	SoldCountPercentage	ABC_Analysis
246.3772635	246.3772635	0	A
495.6922047	495.6922047	0	A
243.333373	243.333373	0	A
200.2031968	182.7596783	4.08E-05	A
384.5214701	355.9977058	8.17E-05	A
464.9083324	464.9083324	8.17E-05	A
658.9775111	658.9775111	8.17E-05	A
263.4470176	263.4470176	8.17E-05	A
283.165274	283.165274	8.17E-05	A

شکل ۱۱: افزودن محاسبات EPQ و EOQ

اهداف تحلیلی و پیاده‌سازی

هدف اصلی برای پیش‌پردازش و تولید یک مدل یادگیرنده و پیش‌بینانه از دادگان موجود و استاندارد که در این مقاله مورد استفاده قرار خواهد گرفت، اعم از شناسایی هدف، نرمال‌سازی، آموزش، ساخت مدل شبکه عصبی LSTM برای کنترل موجودی در انبار و پیدایش بینشی برای شناسایی موجودی‌ها قبلی و موجودی‌های فعال و پیش‌بینی فروش در بازه‌های زمانی پیش رو است. مدل LSTM دو جهت را با استفاده از ویژگی‌های مشخص شده می‌سازد. مدل‌ها با استفاده از دقت آنتروپی متقاطع باینری به‌عنوان معیار ارزیابی آموزش داده و ارزیابی می‌شوند.

ارزیابی مدل پیشنهادی

لازم به ذکر است، داده‌های حاصل از مرحله جهت یادگیری در شبکه عصبی به‌روزرسانی شده و پیش‌پردازش داده‌ها بر اساس محاسبات آنالیز قانون پارتو صورت گرفته است. همچنین، حذف مقادیر بسیار بزرگ با استفاده از تکنیک Mask انجام شده است و نهایتاً نرمال‌سازی ویژگی‌ها با استفاده از کمیت بیشینه و کمینه صورت پذیرفته است. بعد از ساخت مدل LSTM و آموزش این شبکه‌ها با پارامترهای تعیین شده، در راستای هدف مقاله، دقت ۹۳ درصد حاصل گردید.

با پارامترهای تعیین شده LSTM آموزش مدل
model.fit(X_train, y_train, epochs=10, batch_size=32)

```
Epoch 1/10
4910/4910 [=====] - 12s 2ms/step - loss: 0.2022 - accuracy: 0.9340
Epoch 2/10
4910/4910 [=====] - 9s 2ms/step - loss: 0.1770 - accuracy: 0.9350
Epoch 3/10
4910/4910 [=====] - 10s 2ms/step - loss: 0.1720 - accuracy: 0.9355
Epoch 4/10
4910/4910 [=====] - 10s 2ms/step - loss: 0.1709 - accuracy: 0.9356
Epoch 5/10
4910/4910 [=====] - 9s 2ms/step - loss: 0.1701 - accuracy: 0.9359
Epoch 6/10
4910/4910 [=====] - 10s 2ms/step - loss: 0.1696 - accuracy: 0.9361
Epoch 7/10
4910/4910 [=====] - 9s 2ms/step - loss: 0.1691 - accuracy: 0.9362
Epoch 8/10
4910/4910 [=====] - 9s 2ms/step - loss: 0.1686 - accuracy: 0.9361
Epoch 9/10
4910/4910 [=====] - 9s 2ms/step - loss: 0.1684 - accuracy: 0.9358
Epoch 10/10
4910/4910 [=====] - 10s 2ms/step - loss: 0.1682 - accuracy: 0.9361
<keras.src.callbacks.History at 0x7ee9f3e30fa0>
```

شکل ۱۲: صحت دقت مدل LSTM

۴ بحث و نتیجه‌گیری

برای انجام فرایند کنترل موجودی، روش‌ها و مدل‌های مختلفی وجود دارد. هر مدل موجودی دارای رویکرد متفاوت و مخصوص به خود است و بسته به نوع استفاده می‌تواند در هر کسب‌وکاری متفاوت باشد. در این مقاله از سه مدل مهم کنترل موجودی در زنجیره تأمین استفاده نمودیم که هرکدام حسب ویژگی‌های منحصر به خود نسبت به پردازش اولیه داده‌های انبوه در کلان داده اثربخش عمل نمود. اولین مورد، تکنیک کنترل موجودی EOQ بکار برده شد. نه تنها به ما این امکان را داد تا تعداد واحدهای موجودی را که باید برای کاهش هزینه‌ها بر اساس هزینه‌های نگهداری شرکت، همچنین هزینه‌های سفارش و

نرخ تقاضای سفارش شناسایی کنیم، بلکه جهت آموزش به ماشین نیز با دقت مطلوبی اثربخش عمل کرد. دومین مورد، تکنیک EPQ است. در این نوع از مدل‌های کنترل موجودی تعداد محصولات را که باید در یک دسته سفارش داد را تعیین می‌کند، لذا با این هدف که هزینه‌های نگهداری و هزینه‌های راه‌اندازی را کاهش دهد در وزن دهی برای آموزش به ماشین تعریف کردیم. یافته‌ها نشان داد در زمانی که تقاضا برای محصولات در دوره‌های زمانی ثابت باشد، دقت پیش‌بینی را با دقت مطلوبی افزایش خواهد داد. مورد سوم، تکنیک کنترل موجودی آنالیز ABC است، نتیجه مطلوبی که توسط این روش داده‌های موجودی را بر اساس سطوح اهمیت آن‌ها جهت آموزش به ماشین دسته‌بندی کردیم بسیار اثربخش عمل کرد.

از سوی دیگر، یافته‌ها نشان داد، هوشمند سازی زنجیره تأمین در صنایع مختلف، در نتیجه ادغام تکنیک‌های یادگیری عمیق (شبکه‌های حافظه طولانی کوتاه‌مدت LSTM)، با چارچوب‌های سنتی ذکر شده؛ برای پیش‌بینی موجودی انبار دارای عملکرد بالایی به لحاظ سرعت و کاهش هزینه داشته. در این راستا، رویکرد پیشنهاد شده در این مقاله استفاده از شبکه عصبی LSTM است. بردارهای ویژگی برای استخراج الگوهای مصرف وارد شده در شبکه LSTM دقت مدل را بالاتر برد. در این ساختار وزن نوروها و پارامترهای تطبیقی با توجه به الگوهای موجود در ویژگی‌های استخراج شده آموزش داده نیز مطلوب عمل کردند. همچنین فرآیند یادگیری در ساختار LSTM این شبکه را قادر به پیش‌بینی میزان مصرف، در پیش‌آمدهای آتی با دقت ۹۳ درصد نمود. صحت کارآمدی آموزش به شبکه حافظه‌دار در مرحله آموزش؛ شبکه را قادر کرد داده‌های جدید در مورد وضعیت انبار و محصولات را با دقت مطلوبی پردازش کند.

درنهایت، به‌منظور نمایش عملکرد شبکه، نتیجه ارزیابی دقت شبکه عمیق بر روی داده‌ها به شکل عددی و نموداری مقایسه شد. این سیستم نتایج یکنواخت و امیدوارکننده‌ای را برای پیش‌بینی نیازهای حوزه مربوطه به‌منظور کنترل میزان تولید بر اساس تئوری عرضه و تقاضا ارائه کرده است.

الگوی‌های یادگیری نیاز به تعداد زیادی نمونه آموزشی دارد و البته که ممکن است داده‌های قدیمی‌تر در دسترس نباشند، همچنین داده‌های جدید به دلیل مشکلات حفظ حریم خصوصی نتوانند ذخیره شوند یا دفعاتی که سیستم باید در آن بروز رسانی شود، نمی‌تواند از آموزش یک مدل جدید با تمام داده‌ها به اندازه کافی پشتیبانی کند.

از طرفی، با توجه به اینکه استفاده از تکنیک‌های یادگیری عمیق مستلزم وجود داده‌های مناسب جهت یادگیری شبکه عصبی است لذا موضوع سوگیری‌ها نیز یک مشکل عمده برای مدل‌های یادگیری عمیق است. اگر یک مدل بر روی داده‌هایی آموزش ببیند که دارای سوگیری هستند، مدل آن سوگیری‌ها را در پیش‌بینی‌های خود بازتولید می‌کند.

هر مدل استخراج شده منحصر به همان محصول و فاکتورهای تعریف شده حسب چالش‌ها، پتانسیل‌ها و محدودیت‌های موجود است. لذا ممکن است برای تحلیل و ارزیابی محصولات و چالش‌های دیگر مناسب و کارآمد نباشد. از این رو اگر بخواهیم این مدل را برای ارزیابی مصرف محصول دیگر به کار ببریم، بهتر است مدل پیاده‌سازی شده با تغییراتی همراه باشد و پارامترهای تنظیمی شبکه عمیق مجدداً انتخاب شوند؛ که خود موضوع پیشنهادشده‌ای جهت پژوهش‌های آتی خواهد بود.

از آنجاکه یادگیری ماشین پتانسیل نوآوری صنایع و مدل‌های کسب‌وکار آن‌ها را دارد بنابراین، پیاده‌سازی یادگیری ماشین در اینترنت صنعتی اشیا می‌تواند مزایای زیادی از جمله، پیش‌بینی تقاضای مصرف‌کننده در صنایع، تعدیل زنجیره تأمین، نگهداری و پیش‌بینی، کنترل کیفیت و افزایش توان عملیاتی تولید برای صنایع مختلف داشته باشد لذا پژوهشگران می‌توانند جهت کارهای آتی این پژوهش را در صنایع دیگر با استفاده از چارچوب‌های منبع باز برای توسعه مدل یادگیری ماشین در محیط صنعتی انجام دهند.

References

- [1] Amin, Motaharinejad, Mahdi and Memarzadeh, Mahdieh (2017). Big data dictionary. Publication of Dibagaran Art Cultural Institute of Tehran.
- [2] Albayrak Ünal, Ö. Erkeyman, B. & Usanmaz, B. Applications of Artificial Intelligence in Inventory Management: A Systematic Review of the Literature. Arch Computat Methods Eng 30, 2605–2625 (2023).

- [3]Aro-Gordon, S., & Gupte, J. (2016). Review of modern inventory management techniques. *Global Journal of Business & Management*, 1(2), 1-22.
- [4]Duan, W., Cao, Q., Yu, Y., & Levy, S. (2013), Mining online user-generated content: using sentiment analysis technique to study hotel service quality. In 2013 46th Hawaii International Conference on System Sciences (pp. 3119-3128). IEEE.
- [5]Faizi Rad, Mohammad Ali, Pooya, Alireza, Naji Azimi, Zahra, & Amir Haeri, Maryam. (2021). Forecasting demand in university reservation systems with the aim of reducing food waste using neural networks with balanced error function. *Industrial Management*, 13(2), 193-170.
- [6]Fan.x, Lyu.x, Xiao,F, Cia.T, Ding.c(2021), Research on Quick Inventory framwork Based on Deep Neural Network, *Procedia Computer Science*
- [7]Fischer, T., & Krauss, C. (2018). Deep learning with long short-term memory networks for financial market predictions. *European Journal of Operational Research*, 270(2), 654–669.
- [8]García-Barrios (2021) A machine learning based method for managing multiple impulse purchase products: an inventory management approach. *J Eng Sci Technol Rev* 14(1):25–37
- [9]Goodfellow, I., Bengio, Y., and Courville, A. (2016). *Deep Learning*. MIT Press.
- [10]Goodfellow, Ian J et al. (2013). “Challenges in representation learning: A report on three machine learning contests”. In: *International Conference on Neural Information Processing*. Springer.
- [11]Kim, S.W. 2022. An investigation on the direct and indirect effect of supply chain integration on firm performance.The *International Journal of Production Economics*, vol. 119(2). Pp. 328-346.
- [12]Kumar,A. Shankar,R. Aljohani, N. (2020)A big data driven framework for demand-driven forecasting with effects of marketing-mix variables.
<https://doi.org/10.1016/j.indmarman.2019.05.003>.
- [13]Le, J. (2018). The 5 Computer Vision Techniques That Will ChLeCun, Y., Bottou, L., Bengio, Y., & Haffner, P. (1998). Gradient-based learning applied to document recognition. *Proceedings of the IEEE*, 86(11), 2278-2324.
- [14]Lee, J., Kao, H. A., & Yang, S. (2014), Service innovation and smart analytics for industry 4.0 and big data environment, *Procedia Cirp*, 16(1), 3-8.

- [15]Marr, B. (2016), Why everyone must get ready for the 4th industrial revolution /2016/04/05/why everyone must get readyfor4thindustrial revolution/#26be9e2f3f90.
- [16]Mauro, D. A., Greco, M., & Grimaldi, M. (2016). A Formal Definition of Big Data Based on its Essential Features. Library Review, 65(3), 122-135. Retrieved
- [17]Mayer-Schoenberger, V., & Cukier, K. (2013). Big Data: A Revolution That Will TransformHow We Live, Work and Think. John Murray
- [18]Nikolopoulos, K., Punia, S., Schafers, A., Tsinopoulos, C., & Vasilakis, C. (2020). Forecasting and planning during a pandemic: COVID-19 growth rates, supply chain disruptions, and governmental decisions. European Journal of Operational Research. <https://doi.org/10.1016/j.ejor.2020.08.001>.
- [19]punia, S. P. Singh, J. K. Madaan, (2020), A cross-temporal hierarchical framework and deep learningforsupplychainforecasting,Computers&IndustrialEngineering<https://doi.org/10.1016/j.cie.2020.107101>.
- [20]Riyazi, Hamid, Doroudian, Mahmoud, Afshar Najafi, Behrouz. (2022). Inventory control of perishable goods based on shelf space and the effect of non-goods changes along with the commitment of minimum purchase of goods. Industrial Management, 14(1), 168-194.
- [21]Sanaei, Ali (2016). Fourth Industrial Revolution, Isfahan: Academic Jihad, University of Isfahan.
- [22]Scheutz, M., & T, M. (2016). Combining Agent-Based Modeling with Big Data Methods to Support Architectural and Urban Design. Springer International Publishing Switzerland, 18.
- [23]Shao, L., Shum, H. P., & Hospedales, T. (2020). Special Issue on Machine Vision with Deep Learning. International journal of computer vision, 128(4), 771-772.
- [24]Shi, Y. (2014). Big Data: History, Current Status, and Challenges Going Forward. The Bridge, The US National Academy of Engineering, 44(4), 6-11.
- [25]Smith, M. L., Smith, L. N., & Hansen, M. F. (2021). The quiet revolution in machine vision-a state-of-the-art survey paper, including historical review, perspectives, and future directions. Computers in Industry, 130, 103472.
- [26]Sultani, Maryam, Khatami Firouzabadi, Seyed Mohammad Ali, Amiri, Maqsood, & Hajian Heydari, Mojtabi. (2023). A hybrid approach to integrated omnidirectional channel demand forecasting using machine learning - time series clustering with dynamic time convolution algorithm and artificial neural

- networks. Articles in Production and Operations Management, 14(1), 121-140. doi: 10.22108/pom.2023.136202.1485
- [27]Tarabian, Zahra, 2017, Development of the green supply chain model in the conditions of uncertainty of demand for perishable pharmaceutical goods using the meta-innovative method of colonial competition, National Conference of Industrial Management and Engineering of Iran, Isfahan, <https://>
- [28]Toomey, J. W. (2000). Inventory management: principles, concepts and techniques (Vol. 12). Springer Science & Business Media.
- [29]Wamba, S., Akter, S., Edwards, A., Chopin, G., Gnanzou, D., 2015. How 'Big Data' Can Make Big Impact: Findings. rom a Systematic Review and a Longitudinal Case Study. International Journal of Production Economics, In: Press, <http://dx.doi.org/10.1016/j.ijpe.2014.12.031>.
- [30]Wazan, Milad (2022). Deep Learning: From Basics to Building Deep Neural Networks with Python. Tehran: Miyad Andisheh.

COPYRIGHTS

© 2023 by the authors. Licensee Modern Management Engineering Journal. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution 4.0 International (CC BY 4.0) (<http://creativecommons.org/licenses/by/4.0/>).

