

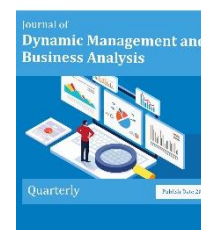


Journal Website

Article history:
Received 30 October 2024
Revised 05 December 2024
Accepted 08 January 2024
Published online 25 January 2025

Dynamic Management and Business Analysis

Volume 4, Issue 1, pp 35-53



E-ISSN: 3041-8933

Estimation of a Financial Distress Prediction Model Based on the Integration of the Support Vector Machine Algorithm and the Least Squares Model

Gholamhasan Taghizad¹, Hossein Panahian²*, Hasan Ghodrati³

1. Phd Student, Department of Accounting, Kashan Branch, Islamic Azad University, Kashan, Iran.

2. Assistant Professor, Department of Accounting, Kashan Branch, Islamic Azad University, Kashan, Iran (Corresponding author).

3. Assistant Professor, Department of Accounting, Kashan Branch, Islamic Azad University, Kashan, Iran.

* Corresponding author email address: h.panahian@iaukashan.ac.ir

Article Info

Article type:

Original Research

How to cite this article:

Taghizad Gh, Panahian H, Ghodrati H. (2025). Estimation of a Financial Distress Prediction Model Based on the Integration of the Support Vector Machine Algorithm and the Least Squares Model. *Dynamic Management and Business Analysis*, 4(1), 35-53.
<https://doi.org/10.61838/dmbaj.173>



© 2024 the author(s). Published by Knowledge Management Scientific Association. This is an open access article under the terms of the Creative Commons Attribution 4.0 International (CC BY 4.0) License.

ABSTRACT

Objective: The objective of this study is to propose a hybrid model based on Partial Least Squares (PLS) and Support Vector Machine (SVM) to predict corporate financial distress and enhance the accuracy and stability of the prediction process.

Methodology: This study utilized a dataset of 120 companies, consisting of 56 bankrupt and 64 non-bankrupt firms, over a two-year period. Initially, financial data were analyzed, and key features were extracted using the Partial Least Squares (PLS) method. The Support Vector Machine (SVM) algorithm was then employed, utilizing a grid search technique with 5-fold cross-validation to optimize model parameters. The performance of the proposed model was compared with traditional methods such as logistic regression and artificial neural networks.

Findings: Empirical results indicated that the hybrid PLS-SVM model achieved an accuracy rate of 87% on the test set, outperforming traditional models and other machine learning techniques. Additionally, the model successfully identified the most relevant financial indicators for predicting financial distress and determined the role of each variable in the prediction process.

Conclusion: Due to its high accuracy, interpretability, and significant stability, the proposed model can serve as an effective tool for financial institutions in risk management, credit approval, and financial planning processes. This study demonstrates that combining machine learning methods can improve financial prediction capabilities.

Keywords: Model, Machine Learning Techniques, Non-linearity, Complex Correlations, Bankruptcy.

EXTENDED ABSTRACT

Introduction

Financial distress prediction has been a critical area of research in financial management, with early efforts dating back to Beaver (1968), who introduced a univariate discriminant analysis model using selected financial ratios (Kollár, 2021; Reyaz et al., 2023). Over time, the field has evolved, incorporating sophisticated statistical and data mining techniques such as logistic regression, decision trees, artificial neural networks (ANN), and support vector machines (SVM). Traditional statistical models, while widely used, have limitations due to assumptions such as linearity, normality, and independence of predictor variables (Hameed et al., 2023; Obaideen, 2024).

Among the modern techniques, SVM has gained prominence due to its ability to manage complex, high-dimensional data while minimizing structural risk rather than empirical risk (Rumbe et al., 2024; Vendittoli et al., 2024). The integration of machine learning methods has led to the development of hybrid models that combine the strengths of different approaches. Partial Least Squares (PLS) is a widely used feature extraction technique that effectively reduces dimensionality and removes multicollinearity, making it suitable for financial distress prediction (Min & Jeong, 2009; Pietruszkiewicz, 2004).

This study aims to develop and evaluate a hybrid financial distress prediction model that integrates PLS with SVM. The proposed model seeks to enhance prediction accuracy and stability by leveraging the strengths of both methods in feature selection and classification. The study addresses the gap in previous research, which often relies on individual techniques rather than hybrid approaches to financial distress prediction (Lee et al., 2002; Malhotra & Malhotra, 2002).

Methodology

The research employed a dataset consisting of 120 firms, including 56 bankrupt and 64 non-bankrupt companies, spanning a two-year period. Initially, financial data were analyzed, and significant features were extracted using the PLS method. PLS facilitated the identification of relevant variables by compressing explanatory variables while preserving their correlation with the target variable.

The SVM algorithm was then applied to classify firms into bankrupt and non-bankrupt categories. The model parameters were optimized using a grid search technique with 5-fold cross-validation. Several kernel functions, including linear, polynomial, and radial basis function (RBF), were tested to determine the optimal classification performance.

The model's performance was evaluated by comparing its accuracy, sensitivity, specificity, and predictive values against traditional methods such as logistic regression and artificial neural networks. The statistical software SPSS23 was used for descriptive analysis, while MATLAB's OSU_SVM toolbox was utilized for SVM classification.

Findings

Empirical results revealed that the hybrid PLS-SVM model achieved an accuracy rate of 87% in the test dataset, surpassing the performance of traditional statistical methods and standalone machine

learning algorithms. The model demonstrated high sensitivity (82%) and specificity (91%), indicating its robustness in distinguishing between bankrupt and non-bankrupt firms.

The study identified key financial indicators that significantly contribute to financial distress prediction. The top features, including profitability ratios, liquidity ratios, and leverage ratios, were determined based on their variable importance in projection (VIP) scores obtained from the PLS analysis.

A comparison with logistic regression and ANN models indicated that the hybrid PLS-SVM model provided superior classification results. The ANN model, despite its high prediction power, suffered from overfitting and required extensive training data, whereas the logistic regression model struggled with nonlinear relationships inherent in financial data.

Additionally, the proposed model demonstrated significant interpretability advantages, as the PLS method allowed for the identification of influential variables, providing valuable insights into financial distress factors.

Discussion and Conclusion

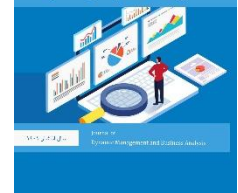
The findings suggest that the integration of PLS and SVM provides a powerful framework for financial distress prediction, offering improved accuracy, stability, and interpretability. The hybrid approach successfully addresses key challenges associated with financial prediction, such as feature selection, multicollinearity, and nonlinearity.

The high prediction accuracy of the proposed model makes it a valuable tool for financial institutions in risk management, credit approval, and strategic financial planning. By utilizing the hybrid model, financial decision-makers can better identify at-risk firms and implement proactive measures to mitigate financial distress.

Compared to traditional models, the PLS-SVM approach demonstrates superior generalization capabilities and reduced sensitivity to overfitting, making it more applicable in real-world financial environments. The study highlights the importance of feature selection and optimization in enhancing the predictive power of financial distress models.

Despite its advantages, the study acknowledges certain limitations, such as the reliance on historical financial data and the need for further validation across different industries and economic conditions. Future research should explore the integration of additional machine learning techniques, such as deep learning, to further improve prediction accuracy and adaptability.

Overall, the study concludes that the proposed hybrid model is a promising advancement in financial distress prediction, offering a reliable and interpretable solution for financial institutions to enhance decision-making processes and risk assessment strategies.



برآورد الگوی پیش بینی درماندگی مالی مبتنی بر تلفیق الگوریتم بردار ماشین و الگوی حداقل مربعات

غلامحسن تقی زاد^۱، حسین پناهیان^{۲*}، حسن قدرتی^۳

۱. دانشجوی دکتری، گروه حسابداری، واحد کاشان، دانشگاه آزاد اسلامی، کاشان، ایران.

۲. استاد تمام، گروه حسابداری، واحد کاشان، دانشگاه آزاد اسلامی، کاشان، ایران (نویسنده مسئول).

۳. استادیار، گروه حسابداری، واحد کاشان، دانشگاه آزاد اسلامی، کاشان، ایران.

*ایمیل نویسنده مسئول: h.panahian@iaukashan.ac.ir

اطلاعات مقاله

چکیده

نوع مقاله

پژوهشی/اصلی

نحوه استناد به این مقاله:

تقی زاد غ، پناهیان ح، قدرتی ح. (۱۴۰۴). برآورد الگوی پیش بینی درماندگی مالی مبتنی بر تلفیق الگوریتم بردار ماشین و حداقل مربعات. مدیریت پویا و تحلیل کسب و کار، ۴(۱)، ۵۳-۳۵.



© ۱۴۰۳ تمامی حقوق انتشار این مقاله متعلق به نویسنده(گان) است. انتشار این مقاله به صورت دسترسی آزاد مطابق با گواهی (CC BY 4.0) صورت گرفته است.

هدف: هدف از این پژوهش، ارائه یک مدل ترکیبی مبتنی بر حداقل مربعات جزئی (PLS) و ماشین بردار پشتیبان (SVM) برای پیش بینی درماندگی مالی شرکت ها و بهبود دقت و پایداری فرآیند پیش بینی است. **روش شناسی:** این پژوهش از یک مجموعه داده شامل ۱۲۰ شرکت، متشکل از ۵۶ شرکت ورشکسته و ۶۴ شرکت غیر ورشکسته، در بازه زمانی دو ساله استفاده کرده است. ابتدا داده های مالی مورد تجزیه و تحلیل قرار گرفتند و ویژگی های مهم با استفاده از روش حداقل مربعات جزئی (PLS) استخراج شدند. سپس از الگوریتم ماشین بردار پشتیبان (SVM) با استفاده از روش جستجوی شبکه ای و اعتبارسنجی متقاطع ۵ بخشی برای تنظیم بهینه پارامترهای مدل استفاده شد. عملکرد مدل پیشنهادی با روش های سنتی مانند رگرسیون لجستیک و شبکه های عصبی مصنوعی مقایسه گردید. **یافته ها:** نتایج تجربی نشان دادند که مدل ترکیبی PLS-SVM با نرخ دقت ۸۷ درصد در مجموعه آزمایشی، عملکرد بهتری نسبت به مدل های سنتی و سایر تکنیک های یادگیری ماشین دارد. همچنین، این مدل توانست مرتبط ترین شاخص های مالی را برای پیش بینی درماندگی مالی شناسایی کرده و تأثیر هر یک از متغیرها را در فرآیند پیش بینی مشخص نماید. **نتیجه گیری:** مدل پیشنهادی به دلیل دقت بالا، تفسیرپذیری مناسب، و پایداری قابل توجه، می تواند به عنوان یک ابزار مؤثر برای مؤسسات مالی در فرآیندهای مدیریت ریسک، تأیید اعتبار، و برنامه ریزی مالی به کار رود. این پژوهش نشان می دهد که ترکیب روش های یادگیری ماشین می تواند به بهبود قابلیت های پیش بینی مالی کمک کند.

کلیدواژه ها: الگو، تکنیک های یادگیری ماشین، غیرخطی بودن، همبستگی های پیچیده، ورشکستگی.

مقدمه

در طی سالیان گذشته و بر اساس تحقیقات و مقالات علمی متعدد، تلاش‌های زیادی برای ایجاد مدل پیش‌بینی شکست کسب و کارها صورت گرفته است. اولین نشریه در سال ۱۹۶۸ منتشر و توسط بیور (۱۹۶۸) که یک مدل تفکیک تک متغیره را با استفاده از نسبت‌های مالی انتخاب شده توسط آزمون طبقه‌بندی دوگانه ایجاد کرده بود، تالیف شد. اکنون، مدل‌های پیش‌بینی ورشکستگی از تکنیک‌های تحلیل آماری و داده‌کاوی برای اصلاح ابزارهای حمایتی و پشتیبان در بهبود تصمیم‌گیری استفاده می‌کنند. علاوه بر تحلیل تشخیصی، روش‌های آماری سنتی شامل رگرسیون، مدل‌های لجستیک، تحلیل عاملی و غیره در کنار تکنیک‌های جدیدتر داده‌کاوی شامل درخت‌های تصمیم‌گیری، شبکه‌های عصبی (NN)، منطق فازی، الگوریتم ژنتیک (GA)، ماشین بردار پشتیبان (SVM) و غیره می‌باشد. کاربردهای آماری، اگرچه در طول زمان افزایش یافتند، اما با مفروضات دقیق آمارهای سنتی مانند خطی بودن، نرمال بودن، استقلال در بین متغیرهای پیش‌بینی‌کننده و شکل عملکردی از قبل موجود مرتبط با متغیرهای معیار و پیش‌بینی‌کننده محدود شدند (Arun & Lakshmi, 2022; Hameed et al., 2023; Kollár, 2021; Obaideen, 2024; Psyridou et al., 2024; Reyaz et al., 2023; Rumbe et al., 2024; Vendittoli et al., 2024).

تاکنون، مدل‌های احتمال خطی و احتمال شرطی چند متغیره، الگوریتم تقسیم‌بندی بازگشتی، هوش مصنوعی، تصمیم‌گیری چند معیاره، برنامه‌ریزی ریاضی برای پشتیبانی از تصمیم‌گیری اعتباری پیشنهاد شده‌اند (Altman et al., 1994; Barniv et al., 1997; Coates & Fant, 1993; Cristianini & Shawe-Taylor, 2000; Curram & Mingers, 1994; Dedene, 2002; Desai et al., 1997; Emel et al., 2003; Fan & Palaniswami, 2000; Fanning & Cogger, 1994; Fletcher & Goss, 1993; Lee et al., 1996; Lee et al., 2002; Lopez & Saidenberg, 2000; Malhotra & Malhotra, 2002; Markham & Ragsdale, 1995; Martin, 1997; Min & Jeong, 2009; Moody, 1992; Pietruszkiewicz, 2004; Sarle, 1995; Smith, 1993). به طور خاص، شبکه‌های عصبی مصنوعی (ANN) اغلب در پیشینه پژوهش‌های قبلی استفاده می‌شوند زیرا قدرت پیش‌بینی آنان بهتر از سایرین شناخته شده است. با این حال، معمولاً گزارش شده است که مدل‌های ANN به مقدار زیادی از داده‌های آموزشی برای تخمین توزیع الگوی ورودی نیاز دارند، و آن‌ها در تعمیم نتایج به دلیل ماهیت بیش‌برازشی خود مشکل دارند. علاوه بر این، برای انتخاب پارامترهای کنترلی از جمله متغیرهای ورودی مربوطه، اندازه لایه پنهان، سرعت یادگیری و حرکت، کاملاً به تجربه یا دانش محققین بستگی دارد تا داده‌ها را پیش‌پردازش کنند (Fanning & Cogger, 1994).

هدف این مقاله استفاده از ماشین‌های بردار پشتیبان (SVM) است، که یک تکنیک یادگیری ماشینی نسبتاً جدیدی است و برای کمک به مشکل پیش‌بینی ورشکستگی و ارائه یک مدل جدید برای بهبود دقت پیش‌بینی استفاده می‌شود. SVM توسط وپنیک (۱۹۹۸) توسعه یافته و به دلیل بسیاری از ویژگی‌های جذاب و عملکرد عالی تعمیم در طیف گسترده‌ای از مسائل و مشکلات محبوبیت پیدا کرده است. همچنین، SVM اصل کمینه‌سازی ریسک ساختاری^۱ (SRM) را در بر می‌گیرد، که نشان می‌دهد نسبت به اصل کمینه‌سازی ریسک تجربی^۲ (ERM) که توسط شبکه‌های عصبی معمولی استفاده می‌شود، برتر و از لحاظ عملکرد بهتر است. SRM یک کران بالای خطای تعمیم را به حداقل می‌رساند در حالی که ERM خطا در داده‌های آموزشی را به حداقل می‌رساند. بنابراین، راه‌حل SVM ممکن است بهینه جهانی باشد در حالی که سایر مدل‌های شبکه عصبی تمایل دارند در یک راه‌حل بهینه محلی قرار بگیرند، و بعید است که با SVM بیش‌برازش اتفاق بیفتد (Akour et al., 2022; Goyal, 2022; Reyaz et al., 2023; Saeed et al., 2018; Santoso et al., 2018; Wang, 2024).

1 structural risk minimization

2 empirical risk minimization

علاوه بر این، با در نظر گرفتن اینکه جستجوی پارامتر بهینه نقش مهمی در ساخت یک مدل پیش‌بینی ورشکستگی با دقت و پایداری بالا ایفا می‌کند، این پژوهش از یک تکنیک جستجوی شبکه‌ای با استفاده از اعتبارسنجی متقاطع ۵ برابری برای یافتن مقادیر پارامتر بهینه تابع هسته SVM استفاده کرده و برای ارزیابی دقت پیش‌بینی SVM، عملکرد آن با تحلیل تشخیصی چندگانه (MDA)، تحلیل رگرسیون لجستیک (Logit) و شبکه‌های عصبی پس انتشار کاملاً متصل سه لایه (BPN) مقایسه گردید.

روش پژوهش

PLS یک روش استخراج ویژگی نظارت شده است. با تجزیه و تحلیل مؤلفه‌های اصلی و سنتز استخراج متغیر، جامع ترین متغیرهای توضیحی که متغیر Y را پیش بینی می‌کردند، استخراج شدند. PLS می‌تواند اطلاعات و نویز سیستم مورد بررسی را جدا کند تا مدل‌های مناسب ایجاد شوند. فشرده سازی ویژگی‌های اطلاعاتی بر اساس PLS، متغیر توضیحی X را فشرده می‌کند و همچنین همبستگی با متغیر پیش بینی شده Y را در نظر می‌گیرد. نتایج آن معانی کاربردی تری خواهند داشت.

PLS اولین متغیر پنهان t_1 را از مجموعه متغیر X استخراج می‌کند. t_1 تا حد امکان اطلاعات تغییرات X را استخراج می‌کند. در همان زمان، اولین متغیر پنهان u_1 را از مجموعه متغیر Y استخراج می‌کند تا بیشترین همبستگی بین t_1 و u_1 را تضمین کند. سپس معادلات رگرسیون با Y و t_1 و معادلات با X و t_1 برقرار می‌شوند. اگر معادلات رگرسیون الزامات دقت را برآورده کنند، الگوریتم خاتمه می‌یابد. در غیر این صورت، دومین متغیر پنهان t_2 از اطلاعات باقیمانده که توسط t_1 از X تفسیر شده است، و u_2 از اطلاعات باقیمانده که توسط t_1 از Y تفسیر شده است، استخراج شد. این فرآیند تکرار شده تا به دقت لازم برسد.

با فرض $X = (x_1, x_2, \dots, x_p)$ که n نمونه از شاخص p بعدی و متغیر پیش بینی Y است، بهینه سازی را می‌توان به صورت زیر

فرموله کرد:

$$\begin{cases} \text{Max} & \text{cov}(Xw_i, Yc_i) \\ \text{s. t.} & w_i'w_i = 1; c_i'c_i = 1 \\ & w_i' \sum_X W_j = 0 \\ & c_i' \sum_Y c_j = 0; \end{cases}$$

که در آن ترکیب خطی $Xw_i t_i$ = دومین متغیر پنهان بود و $\sum_Y = Y'Y$, $\sum_X = X'X$

حل مسئله بهینه سازی فوق (w_i, c_i) را می‌توان در زیر یافت (Lorber et al., 1987):

$$W_i = \begin{cases} \Sigma_{XY} \Sigma_{YX} \text{main eigenvector}, & i = 1 \\ (1 - P_X) \Sigma_{XY} (I - P_Y) \Sigma_{YX} \text{main eigenvector}, & i > 1 \end{cases}$$

$$C_i = \begin{cases} \Sigma_{XY} W_i, & i = 1 \\ (I - P_Y) \Sigma_{YX} W_i, & i > 1 \end{cases}$$

$$\text{where } P_X = (\Sigma_X W) [(\Sigma_X W)^T (\Sigma_X W)]^{-1} (\Sigma_X W)^T,$$

$$P_Y = (\Sigma_Y C) [(\Sigma_Y C)^T (\Sigma_Y C)]^{-1} (\Sigma_Y C)^T, W = (W_{ij}), C = (C_{ij}).$$

از یک طرف، مولفه t_h استخراج شده در محاسبه PLS تا حد ممکن اطلاعات تغییرات X را نشان می‌دهد. از طرف دیگر، برای

توضیح اطلاعات Y تا حد امکان با Y همراه است. برای اندازه گیری توضیحی توان t_h به X و Y، تعریف انواع توان توضیحی تعریف می‌شود

که در آن $r(x_i, x_j)$ ضریب همبستگی بین دو متغیر را نشان می‌دهد.

تعریف ۱ (تبیین تنوع و انباشتگی تبیین تنوع). قدرت تبیین تنوع t_h را به $Rd(y_k; t_h) = r^2(y_k; t_h)$ تعریف شد که همان توضیح تغییرات متغیر بین مولفه t_h و متغیر وابسته y_k است. $Rd(Y; t_h) = \frac{1}{q} \sum_{k=1}^q$ که $Rd(y_k; t_h)$ همان توضیح تغییر متغیر بین مولفه t_h و متغیر وابسته Y است. $Rd(Y; t_1, \dots, t_m) = \sum_{h=1}^m$ که $Rd(Y; t_h)$ همان انباشت توضیح تغییرات بین مولفه $t_1, t_2, \dots, t_m$ و متغیر وابسته Y است. $Rd(y_k; t_1, \dots, t_m) = \sum_{h=1}^m$ که $Rd(y_k; t_h)$ همان انباشت توضیح تغییرات بین مولفه $t_1, t_2, \dots, t_m$ و متغیر وابسته y_k است.

یک وظیفه عملی پیش‌بینی درماتدگی مالی استفاده از کمترین ویژگی برای دستیابی به یک هدف بهینه (مانند حداکثر نرخ شناسایی) است. پژوهش حاضر بر آن است تا تعداد کمی از ویژگی‌ها انتخاب گردد که حجم زیادی از اطلاعات را ارائه می‌دهند. از یک طرف، می‌توان افزونگی و نویز را تا حد امکان حذف کرد. از طرف دیگر، می‌توان هزینه‌های عملیاتی واقعی را به میزان قابل توجهی کاهش داد. انتخاب ویژگی معمولاً به دو مرحله تقسیم می‌شود که در مرحله اول، مقدار مشخصی از ویژگی‌ها را از صدها هزار ویژگی بر اساس روش فیلتر نمایش می‌دهد تا فضای جستجو را کاهش دهد. در مرحله دوم، زیر مجموعه‌ای از ویژگی‌های برجسته را انتخاب می‌کند که بر اساس روش Wrapper شرایط را برآورده می‌کند. نحوه یافتن گروهی از مؤثرترین ویژگی‌ها از ویژگی‌های موجود مشکل اصلی است. پژوهش حاضر یک روش انتخاب ویژگی جهانی جدید بر اساس PLS پیشنهاد کردیم.

برای تجزیه و تحلیل توضیح تغییرات متغیر با X به Y ، پژوهش حاضر به طور کمی تأثیر هر x_i به Y را نشان داد. اهمیت متغیر در طرح ریزی تعریف شده است. تفسیر x_i به Y توسط t_h منتقل می‌شود. اگر قدرت توضیحی t_h به Y قوی باشد و x_i نقش مهمی در ساختار t_h ایفا کند، پس قدرت توضیحی x_i به Y باید قابل توجه در نظر گرفته شود. یعنی برای مولفه‌های t_h روی مقدار بزرگ $Rd(Y; t_h)$ ، اگر مقادیر w_{hi} نیز بالا باشند، تفسیر x_i به Y قوی است.

تعریف ۲ (VIP: اهمیت متغیر در طرح ریزی). با توجه به توضیح تغییرات t_h به Y و تجمع توضیح تغییرات $Rd(Y; t_h)$ و $Rd(Y; t_1, t_2, \dots, t_m)$ از $t_1, t_2, \dots, t_m$ به Y ، در این پژوهش معادله اینگونه تعریف شد:

$$VIP(x_i) = \sqrt{\frac{p}{Rd(Y; t_1, \dots, t_m)} \sum_{h=1}^m Rd(Y; t_h) w_{hi}^2}$$

به عنوان اهمیت متغیر در طرح ریزی x_i به Y که در آن w_{hi} وزن یکم محور w_h است که نشان دهنده سهم حاشیه‌ای x_i که مولفه‌های سازنده t_h است، می‌باشد.

هرچه $VIP(x_i)$ بزرگتر باشد نشان دهنده اهمیت بیشتر در تفسیر x_i به Y است. بنابراین می‌توان از شاخص VIP در انتخاب ویژگی استفاده کرد. پیچیدگی محاسباتی الگوریتم ۳، $O(np)$ می‌باشد که یک چند جمله‌ای خطی p (تعداد ویژگی‌ها) است.

Algorithm 1 $PLS_{VIP}(TrainX_{n \times p}, ClassY_{(n \times g)}, nfac)$

ورودی : $TrainX_{(n \times p)}, ClassY_{(n \times g)}, nfac$

خروجی : vip



شروع	Call $SIMPLS(TrainX_{(n \times p)}, ClassY_{(n \times g)}, nfac)^1$ to get calculated $Rd(X), Rd(Y), W$
بازگشت	$Vip \leftarrow \sqrt{p^* < Rd(Y), W^2 > / Rd(Y)}$ (تعریف 2) vip

پایان

¹ SIMPLS: de Jong proposed simple partial least-squares algorithm, computational complexity is $O(np)$

مراحل دقیق به شرح زیر است:

ابتدا، ماتریس دسته Y به صورت زیر کدگذاری می‌شود:

داده‌های X با ماتریس $n \times p$ نشان داده می‌شوند. ماتریس دسته $Y = (y_{ij})_{n \times g}$ کدگذاری شد که n تعداد نمونه‌ها، p تعداد ویژگی‌ها و g تعداد دسته‌ها است. y_{ij} برچسب دسته است که به این صورت تعریف می‌شود:

$$y_{ij} = I(y_i = j) = \begin{cases} 1 & y_i = j \\ 0 & y_i \neq j \end{cases}; i = 1, \dots, n; j = 1, \dots, g;$$

$y_i = \mathbf{y}_i$ برچسب اصلی است.

ریشه میانگین مجموع مربعات خطای باقیمانده پیش بینی شده (RMPRESS) و سطح معنی داری آن توسط اعتبارسنجی متقاطع ترک یک خروجی^۱ (LOOCV) محاسبه می‌شود. تعداد مؤلفه‌هایی تعیین می‌شوند که سه عامل را ترکیب می‌کنند، تعداد مؤلفه‌هایی با حداقل RMPRESS، حداقل تعداد مؤلفه‌هایی با سطح معنی دار RMPRESS $Prob. > 0.1$ ، و قدرت توضیحی و قدرت تجمعی استخراج شده برای هر متغیر مستقل و متغیر وابسته.

در نهایت، اهمیت متغیر در طرح ریزی (VIP) برای هر ویژگی محاسبه می‌شود:

با استفاده از الگوریتم ۱، واریانس تجمعی توضیح داده شده و بردار وزن مؤلفه‌های $nfac$ قبلی محاسبه می‌شود. و سپس اهمیت متغیر در طرح ریزی (VIP) هر ویژگی به منظور طبقه بندی به دست می‌آید.

بر اساس تئوری یادگیری آماری و با استفاده از اصل به حداقل رساندن ریسک ساختاری، وپنیک و همکاران. ماشین بردار پشتیبان (SVM) را پیشنهاد دادند که یک الگوریتم یادگیری ماشینی است (Cristianini & Shawe-Taylor, 2000). این ماشین می‌تواند توانایی تعمیم خوبی را تحت نمونه‌های محدود به دست آورد. با توجه به مجموعه نمونه به صورت زیر:

$$S_T = \{(X_i, y_i) | X_i \in \mathbb{R}^p, y_i \in \{-1, 1\}, i = 1, 2, \dots, n\}$$

تحت شرایط تفکیک پذیری خطی، SVM طبقه بندی مجموعه نمونه حل شده توسط مسائل برنامه ریزی درجه دوم زیر را تغییر

داد:

$$\min_{w, b} \frac{1}{2} w^T w, \\ s. t. \quad y_i(w^T x_i + b) - 1 \geq 0$$

بردارهای پشتیبان (SVs) نمونه‌هایی هستند که معادلات محدودیت فوق را برقرار می‌کنند. آن‌ها عناصر کلیدی نمونه‌های آموزشی هستند و به مرز تصمیم نزدیک تر هستند. اگر تمام نمونه‌های دیگر برداشته شوند و دوباره آموزش داده شوند، سطح طبقه بندی به دست آمده

یکسان است. این بردارهای پشتیبان هستند که نقش مهمی در حل ابر صفحه^۱ بهینه دارند. طبقه بندی غیر خطی و غیر قابل تفکیک توسط SVM از طریق معرفی تابع هسته و متغیرهای لنگی^۲ حل می شود. پس از معرفی تابع هسته، تصمیم نهایی طبقه بندی SVM به صورت زیر عمل می کند:

$$g(x) = \text{sgn} \left\{ \sum_{i=1}^{SV} y_i \alpha_i^* K(x_i, x) + b^* \right\}$$

جایی که α_i^* عنصر غیر صفر w ، $K(x_i, x)$ تابع هسته، SV تعداد عنصر غیر صفر w است و پیچیدگی محاسباتی آن به تعداد SVها بستگی دارد.

هنگام استفاده از ماشین های بردار پشتیبان برای حل یک مشکل خاص، انتخاب تابع هسته مناسب کلید اصلی است. هیچ راه حل خوبی برای ساخت یک تابع هسته مناسب یا انتخاب توابع هسته برای ساخت SVهای کمتر وجود ندارد. در عمل، متداول ترین توابع هسته موارد زیر هستند:

- تابع هسته چند جمله ای $K(x_i, x) = [(x \cdot x_i + 1)]^d$. طبقه بندی کننده SVM مربوطه چند جمله ای d-order است.
- تابع هسته گاوس $K(x_i, x) = \exp\{-\frac{\|x - x_i\|^2}{2\sigma^2}\}$ ، که در آن σ عرض هسته است. طبقه بندی کننده SVM مربوطه، پایه شعاعی (RBF) است که به آن هسته RBF می گویند.
- تابع هسته سیگموئید $K(x, x_i) = \tanh(v(x \cdot x_i) - \theta)$ ، که در آن v ، θ ثابت هستند. طبقه بندی کننده SVM مربوطه یک پرسپترون دو لایه است که حاوی یک لایه پنهان است. گره لایه پنهان به طور خودکار توسط الگوریتم تعیین می شود و هیچ مشکلی برای حداقل نقطه محلی که محدودیت شبکه عصبی است وجود ندارد.

مجموعه داده های مورد استفاده در این تحقیق مجموعه داده های ورشکستگی درپیشینه پژوهش است (پیتروسکیزویکز، ۲۰۰۴). این مجموعه شامل ۱۲۰ شرکت در بازه دو سال متوالی است. در میان این شرکت ها، ۵۶ شرکت ۲ تا ۵ سال بعد ورشکست شدند. هر شرکت با ۳۰ مشخصه به شرح زیر توصیف می شوند:

- X1 - بدهی های نقدی / جاری.
- X2 - دارایی های کل / نقدی.
- X3 - دارایی های جاری / بدهی های جاری.
- X4 - دارایی های جاری / کل دارایی ها.
- X5 - سرمایه در گردش / کل دارایی ها.
- X6 - سرمایه / فروش در گردش.
- X7 - فروش / موجودی.
- X8 - فروش / مطالبات.
- X9 - سود خالص / کل دارایی ها.
- X10 - سود خالص / دارایی های جاری.

1 Hyper-plane
2 Slack variables



- X11 - سود/فروش خالص.
- X12 - سود/فروش ناخالص.
- X13 - سود/بدهی خالص.
- X14 - سود خالص / حقوق صاحبان سهام.
- X15 - سود خالص / (صاحب صاحبان سهام + بدهی های بلند مدت).
- X16 - فروش / مطالبات.
- X17 - فروش / کل دارایی.
- X18 - فروش / دارایی های جاری.
- X19 - (۳۶۵ * مطالبات)/فروش.
- X20 - فروش/کل دارایی.
- X21 - بدهی ها / درآمد کل.
- X22 - بدهی های جاری/کل درآمد.
- X23 - مطالبات/بدهی ها.
- X24 - سود/فروش خالص.
- X25 - بدهی ها / کل دارایی ها.
- X26 - بدهی ها / حقوق صاحبان سهام.
- X27 - بدهی های بلند مدت / حقوق صاحبان سهام.
- X28 - بدهی های جاری / حقوق صاحبان سهام.
- X29 - درآمد قبل از بهره و مالیات / کل دارایی.
- X30 - دارایی ها/فروش جاری.

یافته ها

این پژوهش از داده های سال اول همه شرکت ها (مجموع ۱۲۰ شرکت) به عنوان مجموعه آموزشی برای آموزش طبقه بندی کننده و از سال دوم داده ها (مجموع ۱۲۰ شرکت) به عنوان مجموعه آزمایشی برای آزمایش طبقه بندی کننده، استفاده کرده است. در مجموعه آموزشی ۵۶ شرکت ورشکسته، و ۶۴ شرکت غیرورشکسته وجود دارند. نحوه توزیع در مجموعه آزمایشی نیز به همین شکل است. نرمال سازی داده ها برای مجموعه داده انجام می شود. برای مجموعه آموزشی: $X_{trn} = (X_1, X_2, \dots, X_{30})$ بیان مجموعه آموزشی به صفر میانگین و انحراف معیار واحد در بین نمونه ها تبدیل می شود. به این معنا که،

$$X_{trn}^* = \frac{X_{trn} - E(X_{trn})}{\sqrt{VarX_{trn}}}$$

سپس مجموعه آزمایشی با توجه به میانگین ها و انحرافات معیار مجموعه آموزشی مربوطه تغییر می کند. به این معنا که،

$$X_{tst}^* = \frac{X_{tst} - E(X_{trn})}{\sqrt{VarX_{trn}}}$$

در این پژوهش ضریب همبستگی پیرسون ویژگی‌های ۳۰ مشخصه را محاسبه کرده و سپس نتایج خوشه بندی شد. نتیجه در شکل ۱ نشان داده شده است. «فاصله» کمتر از ۰.۰۵ به شرح زیر است: X8 و X16، X17 و X20، X9 و X29، X11 و X12، X19 و X30، X26 و X27 با X22 و X21، یعنی ضریب همبستگی پیرسون مربوطه آن‌ها $r > 0.95$.

به ویژه $r(X8, X16) = 1$ ، $r(X17, X20) = 1$ ، $r(X11, X12) = 0.9959$ ، $r(X21, X22) = 0.9993$. این بدان معناست که در تمام ۳۰ مشخصه، افزونگی اطلاعات زیادی وجود دارد. همانطور که می‌دانیم، اطلاعات اضافی ممکن است منجر به کاهش کارایی الگوریتم یا حتی شکست شود.

بنابراین، مشخصه‌های اضافی حذف شدند؛ ابتدا، تمام مشخصه‌ها از نمونه‌های آموزشی بر اساس تعریف روش انتخاب ویژگی PLS رتبه بندی گردید.

در مرحله دوم، مشخصه‌های اضافی را بر اساس ضرایب همبستگی پیرسون حذف شدند. نمودار زیر ضریب همبستگی پیرسون ۳۰ مشخصه و دسته بندی را نشان می‌دهد. در نهایت، مشخصه‌های X11، X16، X19، X20، X22، X26، X28 و X29 حذف می‌شوند.

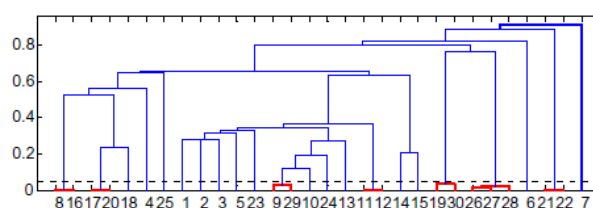
انتخاب ویژگی بر اساس متغیر پیش بینی مهم (VIP) ماشین‌های PLS و SVC اجرا می‌شود. OSU_SVM3.00، جعبه ابزاری از SVMs که از Matlab است، به کار گرفته شده است. تابع هسته خطی، تابع هسته چند جمله‌ای و تابع هسته پایه شعاعی از جعبه ابزار SVMs انتخاب شدند. تنظیمات پارامتر هسته مربوطه در جدول ۲ نشان داده شده است. فرآیند دقیق در زیر ارائه شده است:

مرحله ۱: تمام مشخصه‌های مجموعه آموزشی با روش انتخاب ویژگی که توسط PLS تعریف شده است، مرتب می‌شوند. در این پژوهش بهترین مشخصه‌های k را به عنوان ویژگی‌های مهم انتخاب شد.

مرحله ۲: از ویژگی‌های برجسته k انتخاب شده برای طبقه بندی مجموعه آموزشی با استفاده از ماشین‌های SVC استفاده می‌شود. نتایج در جدول ۲ نشان داده شده است.

شکل ۱

نقشه خوشه بندی ضریب همبستگی پیرسون از ۳۰ ویژگی.



جدول ۱

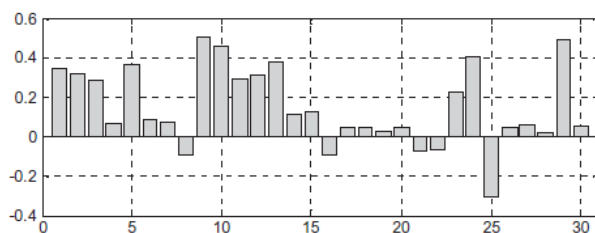
انتخاب ویژگی مبتنی بر PLS

ID	۱	۲	۳	۴	۵	۶	۷	۸	۹	۱۰
مشخصه ها	X1	X2	X3	X4	X5	X6	X7	X8	X9	X10
ID	۱۱	۱۲	۱۳	۱۴	۱۵	۱۶	۱۷	۱۸	۱۹	۲۰
مشخصه ها	X9	X12	X13	X14	X15	X17	X16	X18	X20	X19
ID	۲۱	۲۲	۲۳	۲۴	۲۵	۲۶	۲۷	۲۸	۲۹	۳۰
مشخصه ها	X21	X22	X23	X24	X25	X26	X27	X28	X30	X29



شکل ۲

ضریب همبستگی پیرسون از ۳۰ مشخصه و طبقه بندی.



جدول ۲

انتخاب ویژگی مبتنی بر PLS و مقدار تشخیص طبقه بندی کننده SVM در مجموعه آموزش و آزمایش

عملکرد هسته (پارامتر)	تمامی ۳۰ مشخصه	شماره مشخصه	پس از حذف افزونگی	شماره مشخصه	آزمایشی (۱۲۰)	آموزشی (۱۲۰)
خطی چند جمله ای	۰.۷۸	۱۶ (mfac=۶)	۰.۸۰	۱۰ (mfac=۶)	۰.۸۲	۰.۷۸
(degree=۴)	۰.۹۸	۱۱ (mfac=۳)	۰.۷۹	۱۰ (mfac=۳)	۰.۸۲	۰.۹۹۴
RBF (r=۰.۳, C=۴۰۰)	۰.۹۹	۱۰ (mfac=۳)	۰.۸۳	۷ (mfac=۳)	۰.۸۷	۰.۹۷
طبقه بندی کننده: OSU_SVM۳.۰۰.						

جدول ۳

انتخاب ویژگی مبتنی بر PLS و عملکرد طبقه بندی کننده SVM (RBF)

عملکرد	مجموعه آموزشی	مجموعه آزمایشی
SV	۳۳	۳۵
نرخ صحیح	۰.۹۶۶۷	۰.۸۶۶۷
حساسیت	۰.۹۴۸۳	۰.۸۲۲۶
اختصاصیت	۰.۹۸۳۹	۰.۹۱۰۷
مقدار پیش بینی مثبت	۰.۹۸۲۱	۰.۹۱۰۷
مقدار پیش بینی منفی	۰.۹۸۲۱	۰.۸۲۸۱
طبقه بندی کننده: OSU_SVM۳.۰۰.		

جدول ۴

ویژگی‌های مهم انتخاب شده توسط PLS

ID	پس از حذف افزونگی	نام ملک	مشخصه‌ها
	VIP		
۱	۰.۰۱۰۴	سود خالص / کل دارایی ها	X۹
۲	۰.۰۰۸۳	سرمایه / فروش در گردش	X۶
۳	۰.۰۰۶۵	فروش / کل دارایی	X۱۷
۴	۰.۰۰۵۵	دارایی‌های جاری / بدهی‌های جاری	X۳

۵	۰.۰۰۵۴	X۳۰	فروش جاری / دارایی‌های جاری
۶	۰.۰۰۰۵۰	X۲۳	مطالبات / بدهی‌ها
۷	۰.۰۰۴۵	X۲	دارایی‌ها نقدی / دارایی‌ها کل

مرحله ۳: در مجموعه آموزشی و مجموعه آزمایشی، آزمایشاتی برای محاسبه شاخص‌های مربوطه انجام شد که در جدول ۳ ارائه شده است.

مرحله ۴: مشخصه‌هایی را که حداکثر تشخیص را در مجموعه آزمایشی به دست می‌آورند استخراج شده و زیرمجموعه مشخصه‌های انتخاب شده شرح داده شد. نتایج در جدول ۴ آورده شده است.

پس از حذف افزونگی، ۷ مشخصه از مجموع ۲۲ مشخصه انتخاب شد. نرخ شناسایی SVC برای عملکرد هسته RBF در مجموعه آزمایشی به ۸۷٪ می‌رسد. پس از محاسبه بیشتر هر شاخص، با استفاده از روابط زیر به دست آمد:

$$\text{نرخ صحیح} = \frac{\text{نمونه به‌های درستی طبقه بندی شده}}{\text{نمونه طبقه‌های بندی شده}}$$

$$\text{حساسیت} = \frac{\text{نمونه مثبت‌های به درستی طبقه بندی شده}}{\text{نمونه مثبت‌های واقعی}}$$

مقادیر بیشتر نشان دهنده بهتری از طبقه بندی حساسیت هستند؛

$$\text{اختصاصیت} = \frac{\text{نمونه منفی‌های به درستی طبقه بندی شده}}{\text{نمونه منفی‌های واقعی}}$$

مقادیر بیشتر طبقه بندی دقیق تر و بهتری را نشان می‌دهند،

$$\text{مقدار پیش بینی مثبت} = \frac{\text{نمونه مثبت‌های به درستی طبقه بندی شده}}{\text{نمونه مثبت‌های طبقه بندی شده}}$$

تنگدستی واقعی به عنوان ورشکستگی در نظر گرفته شده است؛

$$\text{مقدار پیش بینی منفی} = \frac{\text{نمونه منفی‌های به درستی طبقه بندی شده}}{\text{نمونه منفی‌های طبقه بندی شده}}$$



غیر ورشکسته واقعی به عنوان غیرورشکسته شناخته شده است. عدد بردار پشتیبان برای مجموعه آموزشی و مجموعه آزمایشی به ترتیب ۳۳ و ۳۵ است. این نشان می‌دهد که تنها ۳۳ و ۳۵ نمونه از مجموع ۱۲۰ نمونه برای دستیابی به شناسایی مناسب در مجموعه‌های آموزشی و آزمایشی مورد نیاز است. این نشان می‌دهد که مدل پیشنهادی پژوهش حاضر توانایی بسیار قوی برای ترویج تعمیم را نشان می‌دهد. در مجموعه آزمایشی، نرخ صحیح تا ۸۷ درصد و نرخ حساسیت به ۸۲ درصد می‌رسد که نشان دهنده حساس بودن مدل طبقه بندی است. اختصاصیت ۹۱ درصد نشان می‌دهد که مدل طبقه بندی دقیق است. شاخص مقدار پیش بینی مثبت نشان می‌دهد که برای شرکت‌های واقعاً ورشکسته نرخ شناسایی تا ۹۱ درصد است. مقدار پیش‌بینی منفی نشان‌دهنده نرخ غیر ورشکستگی است، نرخ شناسایی غیرورشکستگی نیز تا ۸۳ درصد است.

نتایج نشان می‌دهد که پس از حذف مشخصه‌های اضافی، نرخ شناسایی در مجموعه آزمایشی بهبود یافته است، که روش پیشنهادی پژوهش حاضر را تأیید می‌کند که یعنی ابتدا انتخاب ویژگی را هنگام تجزیه و تحلیل داده‌ها با افزودن و نویر انجام گردد. روش پیشنهادی از چندین جنبه عملکرد برتر را نشان می‌دهد: PLS می‌تواند به طور موثر تأثیر هم‌خطی بین مشخصه‌ها را از بین ببرد و دانش کلی را آشکار کند. این یک عملیات جعبه سفید است که بهترین نتایج را ایجاد می‌کند و در عین حال نقش هر ویژگی را در فرآیند پیش بینی ارائه می‌دهد (جدول ۴). مزیت دیگر SVC این است که تابع تصمیم‌گیری نهایی آن تنها توسط تعداد کمی از بردارهای پشتیبانی (SV) تعیین می‌شود. ترکیب PLS و SVC می‌تواند قابلیت مدل‌سازی بسیار خوبی را در مجموعه داده‌های اضافی و پر نویز نشان دهد.

جدول ۵

مجموعه نرخ تشخیص شبکه LVQ از آموزش و آزمایش بر اساس شبکه عصبی رقابتی (همه مشخصه‌ها)

آزمایش	تمامی ۳۰ ویژگی
عملکرد	مجموعه آموزشی
نرخ صحیح	مجموعه آزمایشی
حساسیت	شبکه LVQ بر اساس شبکه عصبی رقابتی (بهینه) تعداد نورون‌ها $k = 19$
اختصاصیت	۰.۶۹۱۷
مقدار پیش بینی مثبت	۰.۷۸۳۳
مقدار پیش بینی منفی	۰.۶۵۰۸
	۰.۷۳۴۴
	۰.۷۳۶۸
	۰.۸۳۹۳
	۰.۷۳۲۱
	۰.۶۵۶۳
	۰.۷۳۴۴

جدول ۶

مجموعه نرخ تشخیص شبکه LVQ از آموزش و آزمایش بر اساس شبکه عصبی رقابتی (پس از حذف افزودن)

آزمایش	پس از حذف ۲۲ افزودن
عملکرد	مجموعه آموزشی
نرخ صحیح	مجموعه آزمایشی
حساسیت	شبکه LVQ در شبکه عصبی رقابتی (بهینه) تعداد نورون‌ها $k = 10$
اختصاصیت	۰.۷۰۰۰
مقدار پیش بینی مثبت	۰.۷۸۳۳
مقدار پیش بینی منفی	۰.۶۶۱۳
	۰.۷۴۱۴
	۰.۷۳۲۱
	۰.۶۷۱۹
	۰.۷۳۴۴

جدول ۷

میزان شناسایی طبقه بندی کننده SVM در مجموعه آموزش و آزمایش

تایع هسته (پارامتر تمامی ۳۰ مشخصه)			پس از حذف ۲۲ افزونگی		
خطی	آموزش (۱۲۰)	آزمایش (۱۲۰)	SV	آموزش (۱۲۰)	آزمایش (۱۲۰)
۰.۸۰	۰.۷۷	۰.۷۸	۳۹.۴۰	۰.۷۸	۴۰.۳۸
چند جمله ای (۲) ۰.۹۵	۰.۷۳	۰.۹۳	۳۱.۳۱	۰.۷۷	۲۷.۳۲
چند جمله ای (۳) ۱.۰۰	۰.۷۷	۱.۰۰	۳۳.۳۲	۰.۷۸	۳۴.۲۷
چند جمله ای (۴) ۱.۰۰	۰.۷۱	۱.۰۰	۳۸.۳۴	۰.۷۲	۳۴.۳۲
RBF ($\gamma = 0.2, C = 20$) ۰.۸۸	۰.۷۹	۰.۸۳	۳۹.۳۹	۰.۷۹	۴۰.۴۰

OSU_SVM3.00 طبقه بندی کننده:

به منظور اعتبارسنجی روش پیشنهادی، مدل های دیگری در پیشینه پژوهش برای مقایسه نتایج اجرا شد.

شبکه عصبی (پرسپترون چند لایه) دارای تحمل خطا، خودسازگاری و توانایی تعمیم قوی است. این شبکه در سال های اخیر، به طور عمده در حوزه های پیش بینی مشکلات مالی و ارزیابی ریسک مالی معرفی شده است. تعداد سلول های عصبی رقابتی می تواند بر عملکرد طبقه بندی شبکه تأثیر بگذارد. برای مجموعه آزمایشی، تعداد سلول های عصبی k هنگامی که سرعت تشخیص آن به حداکثر می رسد، برابر با ۱۹ است. **جدول ۵** نتیجه را هنگامی که تعداد سلول های عصبی k برابر با ۱۹ باشد، نشان می دهد.

به همین ترتیب، پس از حذف ۲۲ مشخصه اضافی، شبکه و تعداد سلول های عصبی $k = 10$ دوباره آموزش دید.

نتایج طبقه بندی و شناسایی در **جدول ۶** ارائه شده است.

در مقایسه با نتایج مدل پیشنهادی، دقت پیش بینی NN بسیار کمتر است. در واقع، می توان نتیجه گرفت که وقتی مشخصه های اضافی حذف شدند، میزان تشخیص مجموعه آموزشی کمی بهبود یافت. با این حال، مقدار تشخیص مجموعه آزمایشی بدون تغییر باقی می ماند. از آنجایی که شبکه عصبی می تواند به هر گونه نقشه برداری غیر خطی از طریق یادگیری نزدیک شود و از محدودیت های یک مدل غیر خطی آزاد است، نیازی به انتخاب مشخصه ها از قبل وجود نداشت. این همچنین نشان می دهد که شبکه عصبی مصنوعی (ANN) قدرت محاسباتی غیر خطی قوی دارد. با این حال، رویکرد شبکه های عصبی (پرسپترون چند لایه) محدودیت هایی دارد. اولاً، اگرچه عملیات جعبه سیاه بهترین نتایج را می دهد، اما نمی تواند چگونگی نقش مشخصه های مربوطه را در فرآیند شناسایی تفسیر کند. ثانیاً، عوامل تجربه ذهنی زیادی مانند اندازه شبکه در انتخاب و آموزش شبکه وجود دارد. در این پژوهش از رایانه ها برای حرکت مداوم شبکه استفاده شده و در نهایت $k = 19$ انتخاب شد ($k = 10$ هنگامی که افزونگی حذف شد). علاوه بر این، شبکه های بسیار پیچیده ممکن است بیش برزش ایجاد کنند، بنابراین توانایی ارتقای تعمیم را از دست می دهند.

مدل پیش بینی مبتنی بر SVC می تواند دقت بالا و عملکرد نمونه کوچک عالی را تضمین کند. طبقه بندی زیر بر اساس SVC به عنوان طبقه بندی کننده، با استفاده از جعبه ابزار OSU_SVM3.00 است. در اینجا هسته خطی، هسته چند جمله ای و تابع هسته RBF انتخاب شد. تنظیمات پارامتر هسته مربوطه و نتایج محاسبات در **جدول ۷** نشان داده شده است. بهترین نتیجه تشخیص SVC بر اساس



عملکرد هسته RFB است. در مجموعه تست، نرخ تشخیص به ۷۹ درصد می‌رسد، که کمی بهتر از پرسپترون چند لایه می‌باشد. اما هنوز هم بسیار کمتر از مدل پیشنهادی پژوهش حاضر است. واضح است که مدل پیشنهادی پژوهش حاضر نتایج بهتری نسبت به مدل NN و مدل SVM به همراه داشت.

بحث و نتیجه‌گیری

در این مطالعه، با تلفیق روش حداقل مربعات جزئی (PLS) و ماشین بردار پشتیبان (SVM) مدلی نوین برای پیش‌بینی درماندگی مالی ارائه شد. نتایج تجربی نشان دادند که این مدل در مقایسه با روش‌های سنتی و سایر مدل‌های یادگیری ماشینی دقت بالاتری دارد. این یافته‌ها با مطالعات گذشته که بیانگر کارایی بالای الگوریتم‌های یادگیری ماشینی در تحلیل مسائل پیچیده مالی هستند، همخوانی دارد (Altman et al., 1994).

روش پیشنهادی توانست شاخص‌های مالی مرتبط را به‌طور مؤثری شناسایی کند و اثرگذاری هر یک از این شاخص‌ها را در پیش‌بینی ورشکستگی تعیین نماید. این ویژگی مدل پیشنهادی نسبت به روش‌های سنتی مانند رگرسیون لجستیک و تحلیل تشخیصی چندگانه، که اغلب به مفروضات سخت‌گیرانه‌ای نظیر نرمال بودن داده‌ها و استقلال متغیرها متکی هستند، یک مزیت رقابتی مهم محسوب می‌شود (Min & Jeong, 2009; Soori et al., 2023).

در مقایسه با مدل‌های شبکه عصبی مصنوعی (ANN)، مدل ترکیبی PLS-SVM عملکرد بهتری نشان داد. اگرچه شبکه‌های عصبی قدرت پیش‌بینی بالایی دارند، اما مشکل بیش‌برازشی و نیاز به داده‌های آموزشی گسترده، باعث کاهش توان تعمیم‌دهی آن‌ها می‌شود (Cho et al., 2024; Mumali & Kałkowska, 2024; Mumali, 2022; et al., 2009). در حالی که مدل پیشنهادی، با استفاده از PLS توانست همبستگی‌های پیچیده میان شاخص‌های مالی را مدیریت کرده و با بهره‌گیری از SVM، پایداری و دقت پیش‌بینی را بهبود بخشد.

همچنین مقایسه عملکرد مدل‌های مختلف نشان داد که حذف افزونگی داده‌ها تأثیر چشمگیری در بهبود عملکرد پیش‌بینی‌کننده‌ها دارد. پس از حذف ویژگی‌های غیرضروری، نرخ دقت مدل پیشنهادی به ۸۷ درصد افزایش یافت که نشان‌دهنده توانایی بالای PLS در استخراج ویژگی‌های معنادار از مجموعه داده‌های پیچیده است (Fanning & Cogger, 1994).

از دیگر مزایای مدل پیشنهادی، توانایی آن در تفسیرپذیری بهتر نسبت به روش‌های دیگر است. روش‌های شبکه عصبی به دلیل ماهیت "جعبه سیاه" بودن، امکان ارائه توضیحات شفاف در خصوص نقش متغیرهای مالی را ندارند (Srinivasan & Kim, 1988; Srinivasan & Ruparel, 1990). اما روش PLS، ضمن کاهش ابعاد داده، به تحلیل نقش هر متغیر در خروجی پیش‌بینی کمک می‌کند. استفاده از مدل ترکیبی PLS-SVM می‌تواند مزایای متعددی را برای مؤسسات مالی به همراه داشته باشد. این مدل می‌تواند در فرایند تأیید اعتبار و رتبه‌بندی اعتباری مشتریان مورد استفاده قرار گیرد. دقت بالای این مدل باعث می‌شود مؤسسات مالی بتوانند با اطمینان بیشتری به تخصیص منابع مالی بپردازند و از ریسک‌های ناشی از اعطای وام به شرکت‌های دارای مشکلات مالی اجتناب کنند (Elmer & Borowski, 1988).

همچنین این مدل در مدیریت سبد وام‌ها و تصمیم‌گیری‌های راهبردی بانکی نقش بسزایی دارد. مدیران مالی می‌توانند از این مدل برای تحلیل پرتفوی خود و کاهش ریسک‌های اعتباری استفاده نمایند (Davis et al., 1992). از سوی دیگر، شرکت‌های سرمایه‌گذاری و مؤسسات بیمه می‌توانند از این مدل به عنوان ابزاری برای مدیریت ریسک و برنامه‌ریزی مالی بهره ببرند.

با وجود مزایای قابل توجه مدل پیشنهادی، برخی محدودیت‌ها نیز در این مطالعه وجود دارد. اولاً، مدل ارائه شده برای یک مجموعه داده خاص طراحی و ارزیابی شده است و ممکن است برای سایر صنایع نیاز به اصلاحاتی داشته باشد. بنابراین، تعمیم‌پذیری مدل باید با داده‌های بیشتر و متنوع‌تری مورد بررسی قرار گیرد (Coates & Fant, 1993).

ثانیاً، انتخاب ویژگی‌های مالی مناسب از طریق PLS گرچه مزایای زیادی دارد، اما ممکن است در محیط‌های پویا که شاخص‌های مالی دائماً در حال تغییر هستند، به‌روزرسانی مدل را ضروری سازد. از این‌رو، استفاده از تکنیک‌های انتخاب ویژگی پویا و سازگار با تغییرات محیطی پیشنهاد می‌شود (Min & Jeong, 2009; Min & Lee, 2005).

نتایج این پژوهش نشان می‌دهد که ترکیب مدل PLS با ماشین بردار پشتیبان می‌تواند به عنوان یک روش کارآمد و دقیق برای پیش‌بینی درماندگی مالی شرکت‌ها مورد استفاده قرار گیرد. این مدل، با ارائه دقت بالا و تفسیرپذیری مناسب، به شرکت‌ها و مؤسسات مالی کمک می‌کند تا تصمیمات بهتری در حوزه مدیریت مالی و اعتباری اتخاذ کنند.

پیشنهاد می‌شود تحقیقات آینده بر توسعه مدل‌های ترکیبی دیگر با استفاده از الگوریتم‌های یادگیری عمیق و تحلیل داده‌های کلان تمرکز کنند. همچنین، بررسی تأثیر عوامل اقتصادی کلان بر عملکرد مدل پیشنهادی می‌تواند دیدگاه‌های جدیدی را برای بهبود قابلیت‌های پیش‌بینی ارائه کند.

در مجموع، مدل پیشنهادی می‌تواند به عنوان یک ابزار تصمیم‌گیری مؤثر برای سازمان‌های مالی در فرآیندهای تحلیل ریسک، تخصیص اعتبار و مدیریت پرتفوی به کار رود و به کاهش خطرات مالی و افزایش سودآوری سازمان‌ها کمک کند.

تعارض منافع

در انجام مطالعه حاضر، هیچ‌گونه تضاد منافی وجود ندارد.

مشارکت نویسندگان

در نگارش این مقاله تمامی نویسندگان نقش یکسانی ایفا کردند.

موازن اخلاقی

در انجام این پژوهش تمامی موازن و اصول اخلاقی رعایت گردیده است.

شفافیت داده‌ها

داده‌ها و مآخذ پژوهش حاضر در صورت درخواست از نویسنده مسئول و ضمن رعایت اصول کپی رایت ارسال خواهد شد.

حامی مالی

این پژوهش حامی مالی نداشته است.



References

- Akour, M., Alenezi, M., & Alsghaier, H. (2022). Software Refactoring Prediction Using SVM and Optimization Algorithms. *Processes*, 10(8), 1611. <https://doi.org/10.3390/pr10081611>
- Altman, E. I., Marco, G., & Varetto, F. (1994). Corporate distress diagnosis comparisons using linear discriminant analysis and neural networks. *Journal of Banking and Finance*, 18(3), 505-529. [https://doi.org/10.1016/0378-4266\(94\)90007-8](https://doi.org/10.1016/0378-4266(94)90007-8)
- Arun, C., & Lakshmi, C. (2022). Genetic algorithm-based oversampling approach to prune the class imbalance issue in software defect prediction. *Soft Computing*, 26(23), 12915-12931. <https://doi.org/10.1007/s00500-021-06112-6>
- Barniv, R., Agarwal, A., & Leach, R. (1997). Predicting the outcome following bankruptcy filing: A three-state classification using neural networks. *International Journal of Intelligent Systems in Accounting, Finance and Management*, 6(3), 177-194. [https://doi.org/10.1002/\(SICI\)1099-1174\(199709\)6:3<177::AID-ISAF134>3.0.CO;2-D](https://doi.org/10.1002/(SICI)1099-1174(199709)6:3<177::AID-ISAF134>3.0.CO;2-D)
- Cho, S., Kim, J., & Bae, J. K. (2009). An integrative model with subject weight based on neural network learning for bankruptcy prediction. *Expert Systems with Applications*, 36(1), 403-410. <https://doi.org/10.1016/j.eswa.2007.09.060>
- Coates, P., & Fant, L. (1993). Recognizing financial distress patterns using a neural network tool. *Financial management*, 22(3), 142-155. <https://doi.org/10.2307/3665934>
- Cristianini, N., & Shawe-Taylor, J. (2000). *An Introduction to Support Vector Machines and Other Kernel-Based Learning Methods*. Cambridge University Press.
- Curram, S. P., & Mingers, J. (1994). Neural networks, decision tree induction and discriminant analysis: An empirical comparison. *Journal of Operational Research Society*, 45(4), 440-450. <https://doi.org/10.1057/jors.1994.62>
- Davis, R. H., Edelman, D. B., & Gammerman, A. J. (1992). Machine learning algorithms for credit-card applications. *IMA Journal of Mathematics Applied in Business and Industry*, 4, 43-51. <https://doi.org/10.1093/imaman/4.1.43>
- Dedene, G. (2002). A comparison of state-of-the-art classification techniques for expert automobile insurance claim fraud detection. *The Journal of Risk and Insurance*, 69(3), 373-421. <https://onlinelibrary.wiley.com/doi/abs/10.1111/1539-6975.00023>
- Desai, V. S., Conway, J. N., & Overstreet, G. A. (1997). Credit scoring models in the credit union environment using neural networks and genetic algorithms. *IMA Journal of Mathematics Applied in Business and Industry*, 8, 324-346. <https://doi.org/10.1093/imaman/8.4.323>
- Elmer, P. J., & Borowski, D. M. (1988). An expert system approach to financial analysis: The case of S&L bankruptcy. *Financial management*, 17(3), 66-76. <https://doi.org/10.2307/3666073>
- Emel, A. B., Oral, M., Reisman, A., & Yolalan, R. (2003). A credit scoring approach for the Commercial Banking Sector. *Socio-Economic Planning Sciences*, 37, 103-123. [https://doi.org/10.1016/S0038-0121\(02\)00044-7](https://doi.org/10.1016/S0038-0121(02)00044-7)
- Fan, A., & Palaniswami, M. (2000). Selecting bankruptcy predictors using a support vector machine approach. *Proceedings of the international joint conference on neural networks*,
- Fanning, K., & Cogger, K. (1994). A comparative analysis of artificial neural networks using financial distress prediction. *International Journal of Intelligent Systems in Accounting, Finance and Management*, 3(3), 241-252. <https://doi.org/10.1002/j.1099-1174.1994.tb00068.x>
- Fletcher, D., & Goss, E. (1993). Forecasting with neural networks and application using bankruptcy data. *Information and Management*, 24, 159-167. [https://doi.org/10.1016/0378-7206\(93\)90064-Z](https://doi.org/10.1016/0378-7206(93)90064-Z)
- Goyal, S. (2022). Genetic evolution-based feature selection for software defect prediction using SVMs. *Journal of Circuits, Systems and Computers*, 31(11), 2250161. <https://doi.org/10.1142/S0218126622501614>
- Hameed, S., Elsheikh, Y., & Azzeh, M. (2023). An optimized case-based software project effort estimation using genetic algorithm. *Information and Software Technology*, 153, 107088. <https://doi.org/10.1016/j.infsof.2022.107088>
- Kollár, A. (2021). Betting Models Using AI: A Review on ANN, SVM, and Markov Chain. <https://doi.org/10.31219/osf.io/mr2v3>
- Lee, K., Han, I., & Kwon, Y. (1996). Hybrid neural networks for bankruptcy predictions. *Decision Support Systems*, 18, 63-72. [https://doi.org/10.1016/0167-9236\(96\)00018-8](https://doi.org/10.1016/0167-9236(96)00018-8)
- Lee, T. S., Chiu, C. C., Lu, C. J., & Chen, I. F. (2002). Credit scoring using hybrid neural discriminant technique. *Expert Systems with Applications*, 23, 245-254. [https://doi.org/10.1016/S0957-4174\(02\)00044-1](https://doi.org/10.1016/S0957-4174(02)00044-1)
- Lopez, J. A., & Saidenberg, M. R. (2000). Evaluating credit risk models. *Journal of Banking and Finance*, 24(1-2), 151-165. [https://doi.org/10.1016/S0378-4266\(99\)00055-2](https://doi.org/10.1016/S0378-4266(99)00055-2)
- Lorber, A., Wangen, L., & Kowalski, B. (1987). A theoretical foundation for the PLS algorithm. *Journal of Chemometrics*, 1, 19-31. <https://doi.org/10.1002/cem.1180010105>
- Malhotra, R., & Malhotra, D. K. (2002). Differentiating between good credits and bad credits using neuro-fuzzy systems. *European Journal of Operational Research*, 136(2), 190-211. [https://doi.org/10.1016/S0377-2217\(01\)00052-2](https://doi.org/10.1016/S0377-2217(01)00052-2)

- Markham, I. S., & Ragsdale, C. T. (1995). Combining neural networks and statistical predictions to solve the classification problem in discriminant analysis. *Decision Sciences*, 26(2), 229-242. <https://doi.org/10.1111/j.1540-5915.1995.tb01427.x>
- Martin, D. (1997). Early warning of bank failure: A logit regression approach. *Journal of Banking and Finance*, 1, 249-276. [https://doi.org/10.1016/0378-4266\(77\)90022-X](https://doi.org/10.1016/0378-4266(77)90022-X)
- Min, J. H., & Jeong, C. (2009). A binary classification method for bankruptcy prediction. *Expert Systems with Applications*, 36, 5256-5263. <https://doi.org/10.1016/j.eswa.2008.06.073>
- Min, J. H., & Lee, Y.-C. (2005). Bankruptcy prediction using support vector machine with optimal choice of kernel function parameters. *Expert Systems with Applications*, 28(4), 603-614. <https://doi.org/10.1016/j.eswa.2004.12.008>
- Moody, J. E. (1992). The effective number of parameters: An analysis of generalization and regularization in nonlinear learning systems. *NIPS*,
- Mumali, F. (2022). Artificial neural network-based decision support systems in manufacturing processes: A systematic literature review. *Computers and Industrial Engineering*, 165, 107964. <https://doi.org/10.1016/j.cie.2022.107964>
- Mumali, F., & Kałowska, J. (2024). Intelligent support in manufacturing process selection based on artificial neural networks, fuzzy logic, and genetic algorithms: Current state and future perspectives. *Computers and Industrial Engineering*, 193, 110272. <https://doi.org/10.1016/j.cie.2024.110272>
- Obaideen, K. (2024). Autonomous Unmanned Systems: Traversing the Bibliometric Terrain of Genetic Algorithm-Based Path Planning. 8. <https://doi.org/10.1117/12.3013834>
- Pietruszkiewicz, W. (2004). *Application of discrete predicting structures in an early warning expert system for financial distress* [Szczecin Technical University]. Szczecin. <https://www.researchgate.net/publication/265382972>
- Psyridou, M., Koponen, T., Tolvanen, A., Aunola, K., Lerkkanen, M.-K., Poikkeus, A.-M., & Torppa, M. (2024). Early prediction of math difficulties with the use of a neural networks model. *Journal of Educational Psychology*, 116(2), 212-232. <https://doi.org/10.1037/edu0000835>
- Reyaz, N., Ahamad, G., Khan, N. J., Naseem, M., & Ali, J. (2023). SVMCTI: Support Vector Machine-Based Cricket Talent Identification Model. <https://doi.org/10.21203/rs.3.rs-2727187/v1>
- Rumbe, G., Hamasha, M., & Mashaqbeh, S. (2024). A comparison of Holts-Winter and Artificial Neural Network approach in forecasting: A case study for tent manufacturing industry. *Results in Engineering*, 21, 101899. <https://doi.org/10.1016/j.rineng.2024.101899>
- Saeed, S., Baber, J., Bakhtyar, M., Ullah, I., Sheikh, N., Dad, I., & Sanjrani, A. A. (2018). Empirical evaluation of SVM for facial expression recognition. *International Journal of Advanced Computer Science and Applications*, 9(11). <https://doi.org/10.14569/ijacsa.2018.091195>
- Santoso, M., Sutjiadi, R., & Lim, R. (2018). Indonesian Stock Prediction Using Support Vector Machine (SVM). *Matec Web of Conferences*, 164, 01031. <https://doi.org/10.1051/mateconf/201816401031>
- Sarle, W. S. (1995). Stopped training and other remedies for overfitting. Proceedings of the twenty-seventh symposium on the interface of computing science and statistics,
- Smith, M. (1993). *Neural networks for statistical modeling*. Van Nostrand Reinhold. <https://archive.org/details/neuralnetworksfo0000smit>
- Soori, M., Arezoo, B., & Dastres, R. (2023). Artificial neural networks in supply chain management, a review. *Journal of Economy and Technology*, 1, 179-196. <https://doi.org/10.1016/j.ject.2023.11.002>
- Srinivasan, V., & Kim, Y. H. (1988). Designing expert financial systems: A case study of corporate credit management. *Financial management*, 5, 32-43. <https://doi.org/10.2307/3666070>
- Srinivasan, V., & Ruparel, B. (1990). CGX: An expert support system for credit granting. *European Journal of Operational Research*, 45, 293-308. [https://doi.org/10.1016/0377-2217\(90\)90194-G](https://doi.org/10.1016/0377-2217(90)90194-G)
- Vendittoli, V., Polini, W., Walter, M. S. J., & Geißelsöder, S. (2024). Introducing Artificial Neural Networks to predict the dimensional and micro-geometrical deviations of additively manufactured parts. *Procedia CIRP*, 129, 181-186. <https://doi.org/10.1016/J.PROCIR.2024.10.032>
- Wang, S. (2024). SVM-based Support Vector Type Recognition Machine for Smart Things in Soccer Training Motion Recognition. *Scalable Computing Practice and Experience*, 25(4), 2519-2531. <https://doi.org/10.12694/scpe.v25i4.2923>