

Risk Assessment Modeling of Urban Railway Investment and Financing Based on Improved SVM Model for Advanced Intelligent Systems

Rupeng Ren, School of Civil Engineering and Architecture, Wuhan University of Technology, China

Jun Fang, School of Civil Engineering and Architecture, Wuhan University of Technology, China

Jun Hu, Ningxia Construction Investment Group Corp., Ltd., China*

Xiaotong Ma, School of Civil Engineering, North Minzu University, China

Xiaoyao Li, Ningxia Construction Investment Group Corp., Ltd., China

ABSTRACT

A risk assessment method for urban railway investment and financing based on an improved SVM model under big data is proposed. First, the inner product in the traditional SVM is replaced by a kernel function to obtain a more accurate non-linear SVM, and a classifier with high classification accuracy is achieved by finding the optimal separating hyperplane. Then, a risk index system is constructed based on the grounded theory combining with intuitionistic fuzzy sets, interval intuitionistic fuzzy sets, weighted averaging operators and the distance measure, and the selection method of assessment indexes is analyzed based on the statistical methods. Finally, the SVM model with fuzzy membership is obtained by fuzzifying the input samples of the SVM based on the given rules of fuzzy membership design. The results show that the maximum relative error between the final test results and the actual value is 0.316%, and the minimum relative error is 0.133% with three different test sets being tested in the proposed method, which can accurately assess the investment.

KEYWORDS

Advanced Intelligent Systems, Big data analysis, Data and Knowledge-Jointly Driven, Mathematical Modeling, Mobile computing environment, SVM, Urban railway

INTRODUCTION

In recent years, with a series of technological breakthroughs such as higher computational intelligence, faster speed, lower energy consumption, and greater comfort, the intelligent subway has become a collection of “intelligent manufacturing,” “intelligent transportation,” and “smart city” technology, comprehensively demonstrating the strength and style of China’s intelligent manufacturing in urban rail transit. With the steady development of China’s economy and society and the acceleration of the

DOI: 10.4018/IJSWIS.331596

*Corresponding Author

This article published as an Open Access article distributed under the terms of the Creative Commons Attribution License (<http://creativecommons.org/licenses/by/4.0/>) which permits unrestricted use, distribution, and production in any medium, provided the author of the original work and original publication source are properly credited.

country's urbanization process, the urban railways have ushered in a golden period of rapid expansion (Tay & Mourad, 2020; Dhaini & Mansour, 2021; Tout et al., 2021). The introduction of competitive market mechanisms for investment and financing of urban railways has led to the emergence of diversified investors, varied financing modes, and complex financing structures (Chehab & Mourad, 2020; Peñalvo et al., 2022). These features have, to some extent, solved the difficulty in financing urban railway projects; but the long construction periods, large investment scale, extensive coverage, and technical complexity of urban railway projects also bring greater risks for rail transit construction (Stergiou et al., 2021; Al-Qerem et al., 2020; Ramaru et al., 2022). Therefore, a correct understanding and analysis of the investment and financing risks of urban railways on the one hand, and an accurate assessment of these projects on the other, are important means of effectively controlling the investment and financing risks of urban railways, as well as achieving the sustainable development of urban railways (Alakbarov, 2022; Lu et al., 2021; Lin, 2021).

Considerable research on urban railway investment and financing risk assessment methods has been conducted, yielding notable results. By identifying the risk factors at all stages of the urban rail transit financing process, a risk evaluation model for urban rail transit project financing has been established based on the extension theory of Liu et al. (2019). However, this method takes each financing risk as an independent factor, ignoring the dynamic correlation among them. Wong et al. (2022) explored the efficacy of fiscal policy in addressing the financing challenges of high-speed rail construction, and conducted a study on the economic and environmental benefits of high-speed rail construction investment for urban development, using the statistics of high-speed rail construction in Chinese cities from 2003 to 2018. However, this study relies too much on the knowledge and experience of experts, which is subjective. Lee et al. (2021) provided a comparative analysis of different assessment criteria and tools used by the private sector, financial institutions, and government contracting agencies to assess the feasibility of urban rail projects and to transfer risks to the party best suited to handle them by optimizing the allocation of risks to minimize the cost. However, this approach only transfers the risks and does not fundamentally offer a solution corresponding to the financing risk assessment. Lv et al. (2020) observed that partners involved in PPPs (Public-Private Partnership) share common interests but come into conflict regarding the value of government subsidies. These authors proposed a method to address this problem by calculating the equitable subsidy ratio favored by all participants, taking into account the uncertain nature of PPPs and the incomplete information used in the decision-making process. However, this method becomes less effective with an excess of sample data, as it is susceptible to interference from redundant information, which reduces learning efficiency.

AlKheder et al. (2022) studied the best way to implement and subsequently manage the new metropolitan line in Kuwait, designing a PPP scheme based on an analysis of the status quo of transport system organization, financing and management schemes, infrastructure projects, and rail service provision. These authors thus developed a risk definition and sharing framework for key elements of the PPP contract. However, the generalization capability of the approach needs to be improved. Aiming to address the issue of the immense financial burden posed by minimum revenue guarantee (MRG) when contract revenue is set considerably higher than actual revenue, Kim et al. (2019) studied how these exercise conditions change the economic value of the MRG, using a case study based on the urban railway project in the Republic of Korea. However, this method does not offer a corresponding strategy for controlling investment and financing risks.

Existing research suggests that findings concerning urban rail transit's investment and financing risks are largely qualitative. There's a notable absence of precise risk measurement and evaluation, making it challenging to lay a solid foundation for risk control. In recent years, a series of reforms have been implemented in the investment and financing system for urban rail transit. The introduction of market competition mechanisms, diversification of investment entities, diversification of financing models, and complexity of financing structures has increased the complexity of risks faced. Various risks and complex relationships exist in the investment and financing of rail transit, and many risk

factors and degrees are difficult to accurately describe. For example, advanced technology, reasonable design of contract terms, clear responsibilities and obligations of participating parties, and high risk levels are all concepts with unclear boundaries, leaving room for ambiguity.

Based on the above analysis, we propose a risk assessment method for urban railway investment and financing based on an improved Support Vector Machine (SVM) model under big data, which will address the problem of low accuracy and large relative error found within the traditional risk assessment methods for urban railway investment and financing. Compared with traditional detection methods, the main contributions of the proposed method are as follows:

- (1) The inner product in the traditional SVM is replaced by a kernel function to obtain a non-linear SVM, which will improve the assessment accuracy of the model.
- (2) The input samples of the tradition SVM are fuzzified to obtain an SVM model with fuzzy membership, which will reduce relative error in the assessment.

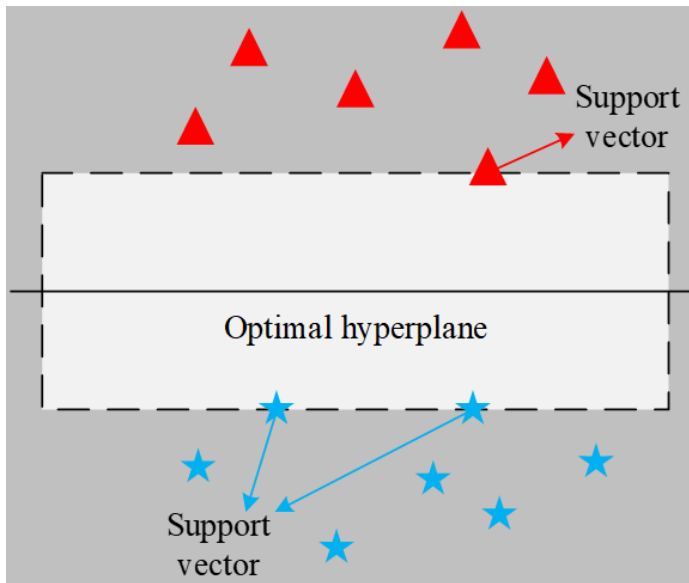
The remaining sections of this paper are arranged as follows: First, we introduce the evaluation methods of our study; second, we demonstrate the construction process of the risk index system; third, we show the multi-factor risk assessment model based on SVM; fourth, we present the experiment itself; and lastly, we provide a conclusion to summarize the study.

ASSESSMENT METHODS

SVM has received extensive attention. It is a classifier based on the principle of minimizing structural risk, which obtains high classification accuracy by finding the optimal separating hyperplanes. The basic principle of the linear separable SVM is shown in Figure 1.

Consider the dataset $A = \{x_k, y_k\}$, where $k = 1, 2, 3, \dots, K$. K is the total number of samples and $x_k \in R^q \subset R$, where x_k is a q dimensional vector. $y_k \in \{-1, 1\}$ is the class label in the binary classification. In the classification, the SVM tries to find the classifier $f(x)$ that can achieve the

Figure 1. Generation process of sentence-level vectors of Word2vec



minimal expected classification error. The linear classifier $f(x)$ is a hyperplane that can be represented as $f(x) = \text{sgn}(\omega^T x + c)$.

The process of finding the optimal classifier $f(x)$ for the SVM is equivalent to optimizing the convex quadratic programming problem shown in Equation (1):

$$\max_{t,c} \frac{1}{2} \|\omega\|^2 + Z \sum_{k=1}^K \gamma_k \quad (1)$$

The constraints of Equation (1) can be expressed as:

$$y_k (< \omega, x_k > + c) \geq 1 - \gamma_k \quad (\gamma_k \geq 0, k = 1, 2, 3, \dots, K) \quad (2)$$

where, Z is the regularization parameter that balances the time complexity against the classification accuracy of the classifier in the dataset A . The above quadratic programming problem can be solved by the dual function. Based on the kernel method, the linear SVM can be converted into a more complex non-linear SVM. Specifically, the kernel function is used to replace the inner product in the above equation (Niu et al., 2020; Zhou et al., 2020). Some typical kernel functions are shown as follows:

$$\text{Lin} : g(x_k, x_l) = x_k^T x_l \quad (3)$$

$$\text{Poly} : g(x_k, x_l) = (x_k^T x_l + 1)^a \quad (4)$$

$$\text{RBF} : g(x_k, x_l) = \exp \left(-\frac{\|x_k - x_l\|_2^2}{2\sigma^2} \right) \quad (5)$$

CONSTRUCTION OF A RISK INDEX SYSTEM

The basic ideas are: 1) The inner product in traditional SVM is replaced by the kernel function, which transforms the linear SVM into a non-linear SVM. 2) A risk index system is constructed based on grounded theory, and the selection method of assessment indexes is analyzed based on the statistical methods. 3) The input samples of SVM are fuzzified to obtain the SVM model with fuzzy membership.

Preliminary Construction of an Index System Based on Grounded Theory

Urban railway projects are inherently uncertain and complex. The dynamic interplay among financing risk assessment indexes is pivotal and should not be overlooked in studying these projects' investment and financing risks (Meng et al., 2022; Du et al., 2021; Guo, 2020). Therefore, it is important to determine the weights of indexes. In order to ensure the scientific accuracy of the weights derived, four relevant definitions are first introduced.

- (1) Intuitionistic fuzzy set. Let A be a non-empty set and there is an intuitionistic fuzzy set $M = \{ \langle x, u_M(x), v_M(x) \rangle | x \in A \}$ on A . The element x in A is about the membership and non-membership in set A and it satisfies the following equations:

$$\begin{cases} 0 \leq u_M(x) \leq 1 \\ 0 \leq v_M(x) \leq 1 \\ 0 \leq u_M(x) + v_M(x) \leq 1 \end{cases} \quad (6)$$

$\alpha_M(x) = 1 - u_M(x) - v_M(x)$ indicates the degree of hesitation and $0 \leq \alpha_M(x) \leq 1$.

- (2) Interval intuitionistic fuzzy set. In intuitionistic fuzzy sets, membership and non-membership are distinctly articulated in uncertainty language, leading to what is termed the interval intuitionistic fuzzy set. The interval intuitionistic fuzzy set can be written as:

$$\tilde{M} = \left\{ \left\langle x, [u_{M1}(x), u_{M0}(x)], [v_{M0}(x), v_{M1}(x)] \right\rangle \mid x \in A \right\} \quad (7)$$

where, $u_{M1}(x_{ij}^k)$ and $u_{M0}(x_{ij}^k)$ denote the upper and lower bounds of membership, respectively. $v_{M1}(x_{ij}^k)$ and $v_{M0}(x_{ij}^k)$ denote the upper and lower bounds of non-membership, respectively. They satisfy the following equation:

$$\begin{cases} 0 \leq u_{M0}(x) \leq u_{M1}(x) \leq 1 \\ 0 \leq v_{M0}(x) \leq v_{M1}(x) \leq 1 \\ u_{M1}(x) + v_{M1}(x) \leq 1 \end{cases} \quad (8)$$

The degree of hesitation can be written as:

$$\alpha_M(x) = [\alpha_{M0}(x), \alpha_{M1}(x)] = [1 - u_{M1}(x) - v_{M1}(x), 1 - u_{M0}(x) - v_{M0}(x)] \quad (9)$$

where $\alpha_{M1}(x)$ and $\alpha_{M0}(x)$ are the upper and lower limits of hesitation, respectively.

- (3) The weighted averaging operator. Let $\tilde{\beta}_i = \left\langle [u_{M0,i}(x), u_{M1,i}(x)] \right\rangle$ be a set of interval intuitionistic fuzzy numbers and $(i = 1, 2, 3, \dots, n)$. The intuitionistic fuzzy weighted averaging operator is ω_{IF} and if $\lambda^n \rightarrow \lambda$, the following equation holds.

$$\begin{aligned} \omega_{\text{IF},\lambda} &= (\tilde{\beta}_1, \tilde{\beta}_2, \tilde{\beta}_3, \dots, \tilde{\beta}_n) = \\ \lambda_1 \tilde{\beta}_1 \oplus \lambda_2 \tilde{\beta}_2 \oplus \lambda_3 \tilde{\beta}_3 \oplus \dots \oplus \lambda_n \tilde{\beta}_n &= \\ \left[\left[1 - \prod_{i=1}^n (1 - u_{M0,i})^{\lambda_i}, 1 - \prod_{i=1}^n (1 - u_{M1,i})^{\lambda_i} \right], \right. \\ \left. \left[\prod_{i=1}^n (v_{M0,i})^{\lambda_i}, \prod_{i=1}^n (v_{M1,i})^{\lambda_i} \right] \right] \end{aligned} \quad (10)$$

where $\lambda = (\lambda_1, \lambda_2, \lambda_3, \dots, \lambda_n)^T$ is the weight vector of $\tilde{\beta}_i$, which satisfies $\lambda_i \in [0, 1]$, and $\sum_i \lambda_i = 1$.

(4) Distance measure. Let two interval intuitionistic fuzzy sets be $\tilde{M}$ and $\tilde{N}$, respectively, which can be written as:

$$\tilde{M} = \left\{ \left\langle x, \left[u_{M1}(x_1), u_{M0}(x_1) \right], \left[v_{M0}(x_1), v_{M1}(x_1) \right] \right\rangle \mid x_1 \in A \right\} \quad (11)$$

$$\tilde{N} = \left\{ \left\langle x, \left[u_{N1}(x_2), u_{N0}(x_2) \right], \left[v_{N0}(x_2), v_{N1}(x_2) \right] \right\rangle \mid x_2 \in A \right\} \quad (12)$$

The distance measure of the interval intuitionistic fuzzy sets is obtained based on the Hemming distance, as shown in Equation (13):

$$\begin{aligned} d(\tilde{M}, \tilde{N}) &= \\ \frac{1}{4} &\left[\left| u_{M1}(x_1) - u_{N1}(x_2) \right| + \left| u_{M0}(x_1) - u_{N0}(x_2) \right| + \right. \\ &\left| v_{M1}(x_1) - v_{N1}(x_2) \right| + \left| v_{M0}(x_1) - v_{N0}(x_2) \right| + \\ &\left. \left| \alpha_{M1}(x_1) - \alpha_{N1}(x_2) \right| + \left| \alpha_{M0}(x_1) - \alpha_{N0}(x_2) \right| \right] \end{aligned} \quad (13)$$

Selection of Indexes Based on Statistical Methods

The construction of the index system involves a comprehensive process of analyzing the results of qualitative and quantitative analyses. The qualitative analysis is based on the preliminary process of risk identification, which mainly consists of the subjective experience of the experts, the statistics and research of previous scholars, and the questionnaires of the actual projects. The quantitative analysis is mainly reflected in the statistical analysis of the scoring results by experts, which are obtained by designed questionnaires (Shih et al., 2019). Through the initial identification of risk factors, the quantitative evaluation of the importance of risk indexes is achieved. And followed by scientific statistical analysis, the original index assessment system is updated to form a second index assessment system. Assuming that the system contains n risk assessment indexes, the basic process of forming the index system is as follows:

(1) Design of the questionnaire. On the basis of establishing a preliminary index system P , a questionnaire survey is conducted on the importance of risk factors of urban railway projects.

The purpose is to develop a quantitative analysis of the index system that accurately reflects the risks of these projects, and which can continuously improve the index assessment system.

- (2) This paper adopts the method of issuing questionnaires. The purpose and meaning of a given questionnaire are first specified, and the basic information of the respondents is collected, such as age, gender, whether they work on transport projects, and years of experience in this type of work. Then actual situations are investigated, and the initial identified n risk factors are introduced into the questionnaire to form a list of risk questionnaires. The experts need to score the risk factors according to their experience, and the score q is in an interval of $[1 - 10]$ with a minimum step of 1. The score ranging from the lowest (1 point) to the highest (10 points) indicates the increasing importance of the risk.
- (3) Distribution and return of questionnaires. Questionnaires are distributed to experts who have been working in urban railways or PPP project management and research for many years. Several sets of data can be obtained through the return of questionnaires (P_n, Q_n) .
- (4) Statistical survey. Statistical survey is carried out when questionnaires are collected after distribution, counting the gender, age, type of employment, and working experience of the participants, and eliminating invalid questionnaires.
- (5) Reliability analysis and factor analysis. The statistical results of the questionnaire survey on project risk factors are ranked using the mean ranking method, where the mean value can be calculated as:

$$A = \frac{\sum(f \times q)}{N}, (1 \leq A \leq 10, 1 \leq q \leq 10) \quad (14)$$

where q indicates the importance of the risk factors as rated by experts. f denotes the frequency of occurrence of each score. N denotes the number of validly interviewed experts. The remaining risk factors are analyzed for reliability and validity by eliminating those with low importance.

- (6) Construct the index system by factor analysis. Factor analysis and dual factor extraction are conducted on the data tested for reliability, obtaining the public factor with high reliability for factor clustering and testing. Then the result of factor clustering is obtained, and the internal information of multiple factors (risk factor $d_1, d_2, d_3, \dots, d_n$) can be expressed by a public factor G . Through naming and interpreting of this public factor, multiple risk factors can be categorized, so as to form the upper layer of risk classification assessment, leading to the index assessment system of risk factors.
- (7) Combined qualitative and quantitative analysis. Despite undergoing rigorous mathematical processing, the returned data might not always produce optimal results. The trends they reflect can be referred to, and the data have practical and theoretical significance. Hence, a combination of qualitative and quantitative analysis is needed to consider the final index system in a comprehensive manner. The final risk index system for investment and financing of urban railway projects is three-tier, as shown in Table 1.

Establishment of Index Weight Model for Investment and Financing Risk Assessment

- (1) By introducing the interval intuitionistic fuzzy numbers, a corresponding five-level linguistic assessment transformation set is established, as shown in Table 2.

Table 1. Three-tier index assessment system for investment and financing risks of urban railways

Investment and financing risks of municipal railway projects			
Type A risk	Type B risk		Type Z risk
Risk factor	Risk factor		Risk factor
$a_1, a_2, \dots, a_n$	$b_1, b_2, \dots, b_n$		$z_1, z_2, \dots, z_n$

- (2) Construct a set of financing risk assessment indexes $S = [s_1, s_2, s_3, \dots, s_n]$ and a set of experts $E = [e_1, e_2, e_3, \dots, e_N]$. Based on the five-level linguistic assessment set, an interval intuitionistic fuzzy scoring set of experts $X = (\tilde{x}_{kl}^m)_{n \times n}$ can be established, where $\tilde{x}_{kl}^m$ indicates the degree of influence that the expert e_m thinks the index s_k has on the index s_l , and $\tilde{x}_{kl}^m = \left\langle [u_0(x_{kl}^m), u_1(x_{kl}^m)], [v_0(x_{kl}^m), v_1(x_{kl}^m)] \right\rangle$.
- (3) Establish a direct fuzzy influence matrix $Y = (\tilde{y}_{kl})_{n \times n}$. Due to the differences in knowledge, experience, and selection preferences among experts, the assessment information is prone to extreme values. And once the decision information has extreme values, it will easily lead to distortion and inaccuracy of the assessment results (Manuylenko et al., 2021). To ensure the accuracy of weights of indexes, the dynamic weight of experts v_{kl}^m is introduced to assemble group decisions of experts, which can be expressed as:

$$f(\tilde{x}_{kl}^m, \tilde{x}_{kl}^s) = 1 - d(\tilde{x}_{kl}^m, \tilde{x}_{kl}^s) \quad (15)$$

$$v_{kl}^m = \frac{\sum_{s=1, s \neq m}^N f(\tilde{x}_{kl}^m, \tilde{x}_{kl}^s)}{\sum_{m=1}^N \sum_{s=1, s \neq m}^N f(\tilde{x}_{kl}^m, \tilde{x}_{kl}^s)} \quad (16)$$

Table 2. A five-level linguistic assessment set transformed based on interval intuitionistic fuzzy numbers

Fuzzy language level	Fuzzy language representation	Interval intuitionistic fuzzy numbers
1	No impact/No risk	$\langle [0.0, 0.2], [0.6, 0.8] \rangle$
2	Weak impact/Weak risk	$\langle [0.2, 0.4], [0.4, 0.6] \rangle$
3	Weak impact/Weak risk	$\langle [0.4, 0.6], [0.2, 0.4] \rangle$
4	High impact/High risk	$\langle [0.6, 0.8], [0.0, 0.2] \rangle$
5	High impact/High risk	$\langle [0.8, 1.0], [0.0, 0.0] \rangle$

$$\tilde{y}_{kl} = \sum_{m=1}^N v_{kl}^m \cdot \tilde{x}_{kl}^m \quad (17)$$

where, $\tilde{y}_{kl} = \left\langle [u_0(y_{kl}), u_1(y_{kl})], [v_0(y_{kl}), v_1(y_{kl})] \right\rangle$ indicates the similarity of expert e_m to expert e_s , and v_{kl}^m indicates the dynamic weight of expert e_m .

- (4) The direct fuzzy influence matrix is normalized to obtain $Z = (\tilde{z}_{kl})_{n \times n}$, which can be calculated as:

$$h_u = \frac{1}{\max_{0 \leq k \leq n} \sum_{l=1}^n u^U(y_{kl})} \quad h_v = \frac{1}{\max_{0 \leq k \leq n} \sum_{l=1}^n v^U(y_{kl})} \quad (18)$$

$$u^o(z_{kl}) = h_u \times u^o(y_{kl}) \quad v^o(z_{kl}) = h_v \times v^o(y_{kl}) \quad (19)$$

where, $\tilde{z}_{kl} = \left\langle [u_0(z_{kl}), u_1(z_{kl})], [v_0(z_{kl}), v_1(z_{kl})] \right\rangle$ and $k, l = 1, 2, 3, \dots, n$, and $o = \{0, 1\}$.

- (5) Calculate the comprehensive fuzzy influence matrix $Q = (\tilde{q}_{kl})_{n \times n}$.

$$[u^o(q_{kl})]_{n \times n} = [u^o(z_{kl})]_{n \times n} \times [I - [u^o(z_{kl})]_{n \times n}]^{-1} \quad (20)$$

$$[v^o(q_{kl})]_{n \times n} = [v^o(z_{kl})]_{n \times n} \times [I - [v^o(z_{kl})]_{n \times n}]^{-1} \quad (21)$$

where, $\tilde{q}_{kl} = \left\langle [u_0(q_{kl}), u_1(q_{kl})], [v_0(q_{kl}), v_1(q_{kl})] \right\rangle$ and $k, l = 1, 2, 3, \dots, n$, and $o = \{0, 1\}$.

- (6) Calculate the comprehensive influence matrix $P = (p_{kl})_{n \times n}$.

$$p_{kl} = \frac{u_0(q_{kl}) + u_1(q_{kl}) - v_0(q_{kl}) - v_1(q_{kl})}{2} \quad (22)$$

- (7) Calculate the degree of influence F_k , the degree of being influenced B_l , the degree of centrality γ_k , and the degree of cause μ_k for each assessment index.

$$\begin{cases} F_k = \sum_{l=1}^n u(q_{kl}), B_l = \sum_{k=1}^n u(q_{kl}) \\ \gamma_k = F_k + B_k, \mu_k = F_k - B_k \end{cases} \quad (23)$$

- (8) Comprehensive analysis of the causality diagram and the actual situation of the project. By comparing the indexes two by two, the judgment matrix is obtained, and the hypermatrix, the weighted hypermatrix, and the limit hypermatrix are calculated in turn to further obtain the weights of each index: $w = (w_1, w_2, w_3, \dots, w_n)$.

- (9) Calculate the mixture weight: $\zeta = w + Q \times w = (I + Q) \times w$, where $\zeta = (\zeta_1, \zeta_2, \zeta_3, \dots, \zeta_n)$, and I is the unit matrix.

MULTI-FACTOR RISK ASSESSMENT BASED ON SVM

SVM With Fuzzy Membership

By fuzzifying the input samples of the SVM, which is selecting an appropriate membership function for fuzzification, each sample (x_k, y_k) is assigned a membership value r_k , and the training set becomes a fuzzy training set, which is shown as:

$$W = \{(x_k, y_k, r_k)\}, k = 1, 2, 3, \dots, n \quad (24)$$

where, $x_k \in R^m$, $y_k \in \{-1, 1\}$, $0 \leq r_k \leq 1$, and when r_k is larger, the higher the confidence degree that the sample belongs to a particular category. The challenge of optimal classification pivots toward identifying the best solution for the objective function, shown in Equation (25) below:

$$\begin{cases} \min H(\lambda) = \frac{\|\lambda\|^2}{2} + C \sum_{k=1}^n r_k \varepsilon_k \\ s.t. \\ y_k (\lambda \cdot x_k + a) + \varepsilon_k \geq 1 \\ \varepsilon_k \geq 0, k = 1, 2, 3, \dots, n \end{cases} \quad (25)$$

As with the standard SVM, C denotes the penalty factor, which can either be a constant or set to a variable factor. The optimization problem can be represented by the Lagrange function, which can be written as:

$$\begin{aligned} L(\lambda, a, \varepsilon, \alpha, \beta) = & \\ & \frac{\|\lambda\|^2}{2} + C \sum_{k=1}^n r_k \varepsilon_k - \sum_{k=1}^n \beta_k \varepsilon_k - \\ & \sum_{k=1}^n \alpha_k [y_k (\lambda \cdot x_k + a) + \varepsilon_k - 1] \end{aligned} \quad (26)$$

where α_k and β_k are Lagrange factors with optimal values obtained at the saddle point, and these two parameters satisfy the conditions shown in Equation (27) below:

$$\begin{cases} \frac{\partial L(\lambda, a, \varepsilon, \alpha, \beta)}{\partial \lambda} = \lambda - \sum_{k=1}^n \alpha_k x_k y_k = 0 \\ \frac{\partial L(\lambda, a, \varepsilon, \alpha, \beta)}{\partial a} = \sum_{k=1}^n \alpha_k y_k = 0 \\ \frac{\partial L(\lambda, a, \varepsilon, \alpha, \beta)}{\partial \varepsilon} = r_k - \alpha_k - \beta_k = 0 \end{cases} \quad (27)$$

At this point, the dual form of the original problem can be expressed as:

$$\begin{cases} \max Q(\alpha) = \sum_{k=1}^n \alpha_k - \frac{1}{2} \sum_{k=1}^n \sum_{l=1}^n \alpha_k \alpha_l y_k y_l K(x_k \cdot x_l) \\ s.t. \\ \sum_{k=1}^n \alpha_k y_k = 0 \\ 0 \leq \alpha_k \leq r_k C, k = 1, 2, 3, \dots, n \end{cases} \quad (28)$$

Solving this quadratic programming problem can yield the optimal solution α_{k0} ; then, the proposed SVM optimal discriminant function can be formulated as:

$$f(x) = \text{sgn}(\lambda_0 \cdot x + a_0), x \in R^m \quad (29)$$

In Equation (29), the following Equation (30) must hold.

$$\begin{cases} \lambda_0 = \sum_{k=1}^n \alpha_{k0} x_k y_k \\ a_0 = y_l - \sum_{k=1}^n \alpha_{k0} y_k (x_k \cdot y_l) \end{cases} \quad (30)$$

The difference between the fuzzy SVM algorithm and the standard SVM algorithm is that the former adds new constraints in the process of solving, and its classification results can be divided into the following four cases:

- (1) When $\alpha_{k0} = 0$, relax variable $\varepsilon_k = 0$, and the corresponding sample x_k is correctly classified.
- (2) When $0 \leq \alpha_{k0} \leq r_k C$ is the relaxation variable $\varepsilon_k = 0$, the corresponding sample x_k is the support vector located near the optimal classification plane, which is the support vector in the general sense.
- (3) When $\alpha_{k0} = r_k C$ is the relaxation variable $\varepsilon_k \neq 0$, the corresponding sample x_k is located near the classification plane, which is easy to be wrongly classified.
- (4) When $\alpha_{k0} > r_k C$, the corresponding sample x_k is wrongly divided. It can be seen that the larger the $r_k C$, the less likely it is that the sample will be wrongly divided. Under opposite conditions, the likelihood that the sample will be wrongly divided increases. Therefore, the design of

membership function is the key to the entire fuzzy algorithm, which requires that the membership function can objectively and accurately reflect the distribution characteristics of samples.

When using the improved fuzzy SVM for credit evaluation, it is important to choose an appropriate membership function to fuzzify the sample data. The chosen membership function should capture the uncertainty of the training samples without excessively complicating the model's computational process. It must objectively and accurately reflect the uncertainty that exists in the samples within the system. The rules for designing the fuzzy membership are as follows:

- (1) Try to make the fuzzy membership of the support vectors close to 1 and the fuzzy membership of the other normal samples greater than or equal to 1. In this way, the selection of fuzzy membership will not alter the support vector's effect on classification.
- (2) As points become increasingly isolated or noisy, the fuzzy membership should gradually decrease to lessen the impact on SVM classification.

Currently, there are three main methods of constructing the membership function:

- (1) Membership based on linear distance. This method is mainly determined according to the importance of the sample in the class, which is the distance from the sample to the class center. Let $\bar{x}$ be the class center; the membership function based on the linear distance can be expressed as:

$$r(x_k) = 1 - \frac{\|x_k - \bar{x}\|}{\max\{\|x_k - \bar{x}\|\}} + \eta \quad (31)$$

where η is pre-assigned and is a very small positive number that ensures all values of $r(x_k)$ are greater than 0.

- (2) Membership based on S-type. By introducing an S-type function, the membership of a sample is regarded as a non-linear relationship between the sample and the center of the class in which it is located. It can be expressed as:

$$r(j_k, u, v, w) = \begin{cases} 1, & j_k \leq u \\ 1 - 2 \left[\frac{j_k - u}{w - u} \right]^2, & u \leq j_k \leq w \\ 2 \left[\frac{j_k - u}{w - u} \right]^2, & v \leq j_k \leq w \\ 0, & j_k \geq w \end{cases} \quad (32)$$

where, j_k is the distance from the sample to the center of the class to which it belongs. u and v are predefined parameters and $v = \frac{u + w}{2}$.

- (3) Membership based on the tightness of samples. Besides the relationship between the sample and the center of the class in which it is located, the connection between samples, which is the tightness of samples, should also be considered. For samples with high tightness and low tightness, they can be calculated in two different ways, as shown in Equation (33):

$$r(x_k) = \begin{cases} \frac{3}{5} \cdot \frac{1 - \frac{j(x_k)}{R}}{1 + \frac{j(x_k)}{R}} + \frac{2}{5}, & j(x_k) < R \\ \frac{2}{5} \cdot \frac{1}{1 + \frac{j(x_k)}{R}}, & j(x_k) \geq R \end{cases} \quad (33)$$

where $j(x_k)$ is the distance from the sample x_k to the center of its smallest bounding sphere.

Assessment Process

The method of combining the rough set and the SVM reasonably employs the advantages of each, using the rough set to filter the index system and inputting it into the SVM to obtain the prediction of feasible outcome. The operational steps are as follows:

- (1) Based on the working process of pattern recognition, a feature space is established, taking the financing risks in urban rail transit projects as the sample.
- (2) The information related to the feature value of financing risks in urban rail transit projects is taken as the input vector of the SVM, and the loss caused by financing risks is taken as the output vector. The fuzzy clustering is used to standardize the data of the established financing risk index system and simplify the indexes in the sample space, so that the essential links between the research objects can be easily deduced and become the core of the sample similarity and classifier design, strengthening the performance of the classifier.
- (3) A polynomial kernel function is selected. Assuming that the training sample set for financing risks is constructed as $\{(x_k, y_k), k = 1, 2, 3, \dots, n\}$, the input for training the SVM is $x_k \in R^h = (x_k^1, x_k^2, x_k^3, \dots, x_k^h)$, which is the index vector of the k th item, while h denotes the dimension of indexes of financing risks, and the output $y_k \in R$ corresponds to the k th item. The optimal parameters are set, starting with the insensitive loss function σ , which is set as $\sigma = 0.0001$.

$$|y - g(x)|_\sigma = \begin{cases} 0, & |y - g(x)| \leq \sigma \\ |y - g(x)| - \sigma, & |y - g(x)| > \sigma \end{cases} \quad (34)$$

Here we assume that the financing risk prediction function is $g(x) = \rho \cdot \Phi(x_k) + b$, where $\rho \in R^h$ is the weight vector, $b \in R$ is the threshold, and $(\cdot)$ is the inner product operation. And by introducing the slack variable s_k, s_{k0} , the optimization problem can be formulated as:

$$\left\{ \begin{array}{l} \min_{\rho, b, s_k, s_{k0}} \frac{|\rho|^2}{2} + C \sum_{k=1}^n (s_k + s_{k0}) \\ s.t. \quad y_k - [\rho \cdot \Phi(x_k) + b] \leq \sigma + s_k \\ \quad [\rho \cdot \Phi(x_k) + b] - y_k \leq \sigma + s_{k0} \\ \quad s_k, s_{k0} \geq 0, k = 1, 2, 3, \dots, n \end{array} \right. \quad (35)$$

By employing the Lagrange transformation, we can express the dual optimization problem as follows:

$$\left\{ \begin{array}{l} \max \left[-\frac{1}{2} \sum_{k=1}^n \sum_{l=1}^n \left(\omega(a_k - a_{k0})(a_l - a_{l0})(x_k, y_k) \right) \right. \\ \quad \left. - \sigma \sum_{k=1}^n (a_k + a_{k0}) \right. \\ \quad \left. + \sum_{k=1}^n y_k (a_k - a_{k0}) \right] \\ s.t. \quad \sum_{k=1}^n y_k (a_k - a_{k0}) = 0 \\ \quad 0 < a_k, a_{k0} < C \end{array} \right. \quad (36)$$

$C \in R$ is the penalty factor. By solving this optimization problem, $\rho = \sum_{k=1}^n (a_k - a_{k0}) \Phi(x_k)$ can be obtained. And the threshold b can be calculated as:

$$b = \frac{1}{N_{SV}} \left\{ \begin{array}{l} \sum_{0 < a_k < C} \left[y_k - \sigma - \sum_{x_l \in SV} (a_k - a_{k0}) x_k x_l \right] + \\ \sum_{0 < a_{k0} < C} \left[y_k + \sigma - \sum_{x_l \in SV} (a_k - a_{k0}) x_k x_l \right] \end{array} \right. \quad (37)$$

Finally, the optimal prediction function can be obtained, as shown in Equation (38):

$$f_0(x) = \sum_{x_i \in SV} (a_k - a_{k0}) \omega_0(x_k, x) + b_0 \quad (38)$$

- (4) Verification of the classification model for the financing risks in urban rail transit projects. Usually, a number of samples are selected from inside and outside the training samples, forming a test set to verify the effectiveness of the classification model. To ensure the stability of the classification model, the test is repeated several times.

EMPIRICAL APPLICATIONS

The experiment was run on the Win10 operating system, using 16GB of memory, and GPU (NVIDIA GeForce RTX 2070) of the Intel Core i7 processor. The algorithm uses the PyCharm editing tool and Python programming language. Based on the proposed model, this paper selects 10 cities comprising

Beijing, Tianjin, Shanghai, Nanjing, Suzhou, Shenyang, Changchun, Guangzhou, Shenzhen, and Chongqing, where urban rail transit has been built or is under construction, as examples for analysis and verification. Quantitative assessment is conducted for the qualitative indexes of financing risks in urban rail transit projects so as to propose reasonable and scientific measures, which will effectively ensure the smooth implementation of the project.

Step 1: Expert scoring. A team of experts engaged in urban rail transit project financing research is set up to score the samples in the financing risk index system, based on their experience. The scoring criteria are: excellent 1–20, good 20–40, moderate 40–60, average 60–80, and poor 80–100 (0 excellent, 1 good, 2 moderate, 3 average, 4 poor). The results are shown in Table 3.

In Table 3, A1 to A7 denote the credit risk, construction risk, operational risk, market risk, financial risk, political risk, and environmental risk respectively. D denotes financing risk. In order to facilitate the subsequent data processing, the qualitative data in Table 3 are quantified and the results are shown in Table 4.

Step 2: Attribute simplification. Condition attributes are eliminated one by one based on the rough set theory, and it is verified whether the elimination introduces new contradiction rules, so as to determine the correctness of attribute simplification. For example, if A6 is eliminated and the remaining attributes are still consistent, then A6 can be eliminated. It is verified that A1, A2, A3, and A4 are the attributes that cannot be eliminated.

Normalize the data by applying the scale conversion method shown in Equation (39):

$$D = \frac{D_0 - D_{0\min}}{D_{0\max} - D_{0\min}} \quad (39)$$

where D_0 is the original data, $D_{0\min}$ is the minimum value of the original data, $D_{0\max}$ is the maximum value of the original data, and D is the processed data.

The final results are shown in Table 5.

Kernel function selection and operation. This paper selects the polynomial kernel function $U(x, x_k) = [(x \cdot x_k) + 1]^h$, setting the optimal parameter $\sigma = 0.0001$ and using $C = 100$ to calculate the coefficient $b_0 = 0.002154$. The data of the first seven indexes in Table 4 are used as training samples, and the corresponding true values of financing risk are used as the training target. The data

Table 3. Assessment data of indexes of financing risks in urban rail transit projects

Item Number	A1	A2	A3	A4	A5	A6	A7	D
1	4.67	16.27	21.96	12.09	3.68	30.66	78.71	13.83
2	34.91	66.73	26.91	28.53	68.11	60.31	67.28	62.66
3	12.59	81.97	29.68	34.37	18.79	36.2	89.57	76.27
4	31.32	63.51	30.96	36.35	22.47	48.54	73.21	52.59
5	9.48	59.47	25.18	24.83	10.21	21.9	20.5	70.28
6	8.25	45.02	20.03	29.59	8.74	13.43	19.61	64.82
7	7.05	50.78	20.34	27.68	8.54	5.65	17.28	74.16
8	13.92	37.56	21.36	37.7	41.33	28.23	36.23	81.44
9	11.91	45.89	21.96	40.81	12.52	22.74	80.31	66.51
10	8.64	17.87	31.61	31.14	4.09	15.44	14.38	18.54

Table 4. Quantification of qualitative data of indexes

Item Number	A1	A2	A3	A4	A5	A6	A7	D
1	0	0	0	0	0	2	4	0
2	1	3	0	0	0	1	3	0
3	0	0	1	0	0	2	2	1
4	1	2	1	1	3	1	1	0
5	0	0	1	1	2	0	0	2
6	1	3	0	1	0	0	0	0
7	0	0	0	1	0	0	0	2
8	0	1	0	1	1	0	4	0
9	0	0	1	1	0	2	3	3
10	0	0	1	1	0	0	2	2

Table 5. Normalized data of financing risk indexes in urban rail transit projects

Item Number	A1	A2	A3	A4
1	0.0000	0.0071	0.1915	0.0000
2	1.3645	0.9841	0.6751	0.5561
3	0.3118	1.1678	0.9457	0.7537
4	0.7857	0.7469	1.0708	0.8206
5	0.1984	0.8359	0.5061	0.4310
6	0.0528	0.5768	0.0029	0.5920
7	0.0502	0.6408	0.0000	0.5274
8	0.2998	0.4030	0.0977	0.8663
9	0.2572	0.5447	0.1915	0.9716
10	0.1180	0.0000	1.1342	0.6444

of last three indexes are used as test samples, and the corresponding true values are used as the training target. The predicted values of the test samples are given by the SVM, with the actual values as the true values. And the following equation is used as the criterion for evaluating the prediction performance of the SVM:

$$\xi = sqr\left(\frac{1}{m} \sum_{k=1}^m (\tilde{x} - x_k)\right) \quad (40)$$

where ξ is the root mean square error (RMSE), and $\tilde{x}$ and x_k are the index values of the SVM and the k th item, respectively. m is the number of test samples. If ξ is smaller, the prediction performance of the SVM is better and the prediction results are more reliable; if ξ is larger, the prediction performance of the SVM is worse. The results are shown in Table 6.

As can be seen from Table 6, the test results, which are obtained by the proposed risk assessment model for urban railway investment and financing based on the improved SVM model under big

Table 6. Test results and relative error

Test samples	Actual value	Test results	Relative error
8	0.0692	0.0694	0.289%
9	0.1503	0.1505	0.133%
10	0.0633	0.0635	0.316%

data, are very close to the actual values for three different sets of test samples, with the maximum relative error between the test results and the actual values being 0.316% and the minimum relative error being 0.133%, indicating that the model can achieve an accurate assessment of urban railway investment and financing risks. This is due to the replacement of the inner product in the traditional SVM with a kernel function, thus converting the linear SVM into a more accurate non-linear SVM, which greatly improves the accuracy of the risk assessment. In addition, the input samples of the non-linear SVM are fuzzified by selecting the appropriate membership function, which reduces the relative error of risk assessment within the model.

CONCLUSION

In this paper, the sophistication of the knowledge- and data-driven intelligent modeling method is enhanced, a risk assessment method for urban railway investment and financing based on an improved SVM model under big data is proposed, and practical verification of the proposed method is carried out using three test sets. By using the proposed model, quantitative evaluation of indexes of financing risks in urban rail transit projects is conducted in 10 cities, where urban rail transit has been built or is under construction. The results show that the non-linear SVM, obtained by replacing the inner product in the traditional SVM with a kernel function, can improve the model's assessment accuracy. The construction of a reliable risk index system and the reasonable selection of indexes can simplify the process of risk assessment quantification. The SVM model with fuzzy membership achieved by fuzzifying the input samples of the traditional SVM can effectively reduce the relative error of assessment. Future research will continue to delve into accurately identifying the various risk types associated with urban rail transit project investment and financing, and will also probe deeper into their underlying causes. In addition, the relationship between various risk factors and effective control methods will continue to be explored.

AUTHOR NOTE

Funding for this research was provided by the Foundation of Ningxia, China (Grant No. 2022AAC03266); the Outstanding Young Teachers Training Foundation of Ningxia, China (Grant No. NGY2020054); the West Light Foundation of the Chinese Academy of Sciences (Grant No.XAB2021YW14); and the Ningxia Key Research and Development Program, China (Grant No. 2021BEG03022). The authors declare there is no conflict of interest.

REFERENCES

- Al-Qerem, A., Alauthman, M., Almomani, A., & Gupta, B. B. (2020). IoT transaction processing through cooperative concurrency control on fog–cloud computing environment. *Soft Computing*, 24(8), 5695–5711. doi:10.1007/s00500-019-04220-y
- Alakbarov, R. (2022). An optimization model for task scheduling in mobile cloud computing. [IJCAC]. *International Journal of Cloud Applications and Computing*, 12(1), 1–17. doi:10.4018/IJCAC.297102
- AlKheder, S., Abdullah, W., & Al Sayegh, H. (2022). A socio-economic study for establishing an environment-friendly metro in Kuwait. *International Journal of Social Economics*, 49(5), 685–709. doi:10.1108/IJSE-04-2021-0210
- Chehab, M., & Mourad, A. (2020). LP-SBA-XACML: Lightweight semantics based scheme enabling intelligent behavior-aware privacy for IoT. *IEEE Transactions on Dependable and Secure Computing*, 19(1), 161–175. doi:10.1109/TDSC.2020.2999866
- Dhaini, M., & Mansour, N. (2021). Squirrel search algorithm for portfolio optimization. *Expert Systems with Applications*, 178, 114968. doi:10.1016/j.eswa.2021.114968
- Du, G. S., Liu, Z. X., & Lu, H. F. (2021). Application of innovative risk early warning mode under big data technology in internet credit financial risk assessment. *Journal of Computational and Applied Mathematics*, 386(5), 324–331. doi:10.1016/j.cam.2020.113260
- Guo, Y. P. (2020). Credit risk assessment of P2P lending platform towards big data based on BP neural network. *Journal of Visual Communication and Image Representation*, 71(13), 92–100. doi:10.1016/j.jvcir.2019.102730
- Kim, K., Cho, H., & Yook, D. (2019). Financing for a sustainable PPP development: Valuation of the contractual rights under exercise conditions for an urban railway PPP project in Korea. *Sustainability (Basel)*, 11(6), 65–72. doi:10.3390/su11061573
- Lee, H., Yoo, J., Kim, D., Lee, K.-C., Kang, S. W., & Moon, D. S. (2021). Study on risk based value for money assessment on urban rail project. *Journal of Korean Society for Urban Railway*, 9(4), 1117–1129. doi:10.24284/JKOSUR.2021.12.9.4.1117
- Lin, M. (2021). Innovative risk early warning model under data mining approach in risk assessment of internet credit finance. *Computational Economics*, 59(4), 1443–1464. doi:10.1007/s10614-021-10180-z
- Liu, W., Deng, S., & Xu, Z. (2019). The dynamic evaluation of urban rail transit project financing risk. *Systems Engineering*, 34(12), 153–158.
- Lu, J., Shen, J., Vijayakumar, P., & Gupta, B. B. (2021). Blockchain-based secure data storage protocol for sensors in the industrial Internet of Things. *IEEE Transactions on Industrial Informatics*, 18(8), 5422–5431. doi:10.1109/TII.2021.3112601
- Lv, J. N., Zhang, Y. Y., & Zhou, W. (2020). Alternative model to determine the optimal government subsidies in construction stage of PPP rail transit projects under dynamic uncertainties. *Mathematical Problems in Engineering*, 2020(3), 264–272. doi:10.1155/2020/3928463
- Manuylenko, V., Borlakova, A. I., Milenkov, A. V., Bigday, O. B., Drannikova, E. A., & Lisitskaya, T. (2021). Development and validation of a model for assessing potential strategic innovation risk in banks based on data mining-Monte-Carlo in the “Open innovation” system. *Risks*, 9(6), 283–290. doi:10.3390/risks9060118
- Meng, T., Li, Q., Dong, Z., & Zhao, F. F. (2022). Research on the risk of social stability of enterprise credit supervision mechanism based on big data. *Journal of Organizational and End User Computing*, 34(3), 56–62.
- Niu, A. W., Cai, B. Q., & Cai, S. S. (2020). Big data analytics for complex credit risk assessment of network lending based on SMOTE algorithm. *Complexity*, 32(8), 154–162. doi:10.1155/2020/8563030
- Peñalvo, F. J. G., Sharma, A., Chhabra, A., Singh, S. K., Kumar, S., Arya, V., & Gaurav, A. (2022). Mobile cloud computing and sustainable development: Opportunities, challenges, and future directions. [IJCAC]. *International Journal of Cloud Applications and Computing*, 12(1), 1–20. doi:10.4018/IJCAC.312583

- Ramaru, E., Garg, L., & Chakraborty, C. (2022). A hybrid cloud enterprise strategic management system. [IJCAC]. *International Journal of Cloud Applications and Computing*, 12(1), 1–18. doi:10.4018/IJCAC.297091
- Shih, D. H., Hsu, H. L., & Shih, P. Y. (2019). A study of early warning system in volume burst risk assessment of stock with big data platform. In *4th IEEE International Conference on Cloud Computing and Big Data Analysis (ICCCBDA '19)*. IEEE. doi:10.1109/ICCCBDA.2019.8725738
- Stergiou, C. L., Psannis, K. E., & Gupta, B. B. (2021). InFeMo: Flexible big data management through a federated cloud system. *ACM Transactions on Internet Technology*, 22(2), 1–22. doi:10.1145/3426972
- Tay, B., & Mourad, A. (2020). Intelligent performance-aware adaptation of control policies for optimizing banking teller process using machine learning. *IEEE Access : Practical Innovations, Open Solutions*, 8, 153403–153412. doi:10.1109/ACCESS.2020.3015616
- Tout, H., Mourad, A., Kara, N., & Talhi, C. (2021). Multi-persona mobility: Joint cost-effective and resource-aware mobile-edge computation offloading. *IEEE/ACM Transactions on Networking*, 29(3), 1408–1421. doi:10.1109/TNET.2021.3066558
- Wong, Z., Chen, A., Shen, C. R., Wu, D. L., & Hondroyiannis, J. (2022). Fiscal policy and the development of green transportation infrastructure: The case of China's high-speed railways. *Economic Change and Restructuring*, 21(6), 16–24. doi:10.1007/s10644-021-09381-1
- Zhou, H. J., Sun, G., Fu, S., Liu, J., Zhou, X. X., & Zhou, J. Y. (2020). A big data mining approach of PSO-based BP neural network for financial risk management with IoT. *IEEE Access : Practical Innovations, Open Solutions*, 7(13), 154035–154043.