

Research on Supply Chain Toughness Evaluation and Optimization Based on Integrated Learning Algorithm

Xuanshuang Wang¹, Quanpeng Chen², , Jia Chen¹

¹ Southwest Jiaotong University Hope College, Chengdu, Sichuan, 610400, China

² Sichuan Vocational College of Finance and Economics, Chengdu, Sichuan, 610101, China

ABSTRACT

Consolidating and improving supply chain resilience and maintaining supply chain stability and security is an important foundation for promoting the realization of high-quality development. After initially selecting supply chain resilience related indicators, the research is screened and downgraded through factor analysis to establish a supply chain resilience evaluation index system. Subsequently, based on the model integration framework, the supply chain toughness evaluation model with improved Stacking integration model is constructed on the basis of a single machine learning algorithm and an integrated learning algorithm, and the model parameters are adjusted and optimized through the learning curve to achieve the optimal evaluation effect and compared with the existing model. The results show that the Stacking supply chain toughness evaluation model constructed in this paper has a relative error of 23% or less in 3685 enterprise samples and accounts for 98.78%. It shows that the Stacking integrated model established in this paper has good prediction effect and high accuracy, which has certain value and significance to the research of supply chain toughness prediction, and can provide scientific reference basis for enterprises.

Keywords: machine learning algorithm, Stacking integration model, factor analysis method, supply chain toughness evaluation

1. Introduction

In recent years, major shifts in the international situation have led to great changes in the external and internal environments, and more and more industrial supply chains have been impacted by major emergencies from different dimensions [1-3]. It is an urgent task for Chinese enterprises to evaluate and optimize supply chain resilience in a truly practical and effective way to better cope with more possible shocks in the future [4-6].

 Corresponding author.

E-mail address: xueshunihao@163.com (Q. Chen).

Received 01 February 2024; Revised 20 April 2024; Accepted 10 November 2024; Published Online 16 April 2025.

DOI: [10.61091/jcmcc127b-298](https://doi.org/10.61091/jcmcc127b-298)

© 2025 The Author(s). Published by Combinatorial Press. This is an open access article under the CC BY license (<https://creativecommons.org/licenses/by/4.0/>).

The so-called supply chain resilience is the ability of a supply chain to withstand the impact of major emergencies and recover quickly after being hit. The stronger the supply chain resilience, the smaller the negative impact of major shocks (e.g., extreme weather, material supply cutoffs, production and logistics disruptions, demand surges, etc.), and the more robust the supply chain is [7-10]. A supply chain with strong resilience can help enterprises better cope with various risks such as market fluctuations, natural disasters, and supplier disruptions, thus guaranteeing customer satisfaction and enhancing enterprise competitiveness [11-13]. The evaluation of supply chain resilience can be assessed by several indicators, commonly including indicator evaluation (e.g., time resilience, inventory resilience, supplier resilience demand resilience), risk evaluation (e.g., natural disasters, quality problems), and capacity evaluation (enterprise's purchasing ability) [14-17]. Supply chain reliability and resilience are critical to business success. By establishing management approaches such as information sharing and partnerships, diversification and flexibility, risk management and business continuity plans, and continuous improvement and innovation, firms can increase supply chain reliability and resilience, reduce risk and keep their business on track [18-21]. These approaches are applicable not only to large enterprises but also to small and medium-sized enterprises (SMEs), helping them to gain a competitive advantage in a highly competitive market.

Literature [22] takes the coal industry as an example, constructs the evaluation index system of industrial coal industry chain and supply chain toughness, and proposes a comprehensive method to evaluate the toughness level of coal industry chain and supply chain, and puts forward relevant suggestions based on the evaluation results to improve the toughness level of the supply chain. Literature [23] aims to improve the toughness of the coal-to-liquid (CTL) industry chain and supply chain, identifies the factors affecting toughness in terms of absorption and adaptation, and reveals the influencing factors, such as the degree of localization of key technologies, by constructing an evaluation index system, which points out the direction for enhancing the toughness of the CTL industry and supply chain. Literature [24] combines supply chain (SC) and Industry 4.0 key enabling technologies, realizes the creation of visibility and sustainability in SC, facilitates intelligent decision-making, and helps companies develop intelligent systems to overcome risks and improve sustainability. Literature [25] proposes a unified framework for evaluating supply chain reliability and resilience, which can effectively capture risks involving supply, demand, the firm itself, and the external environment, and is conducive to promoting sound business development.

The above study elaborates that supply chain resilience evaluation is important for enterprises to cope with risks and development by constructing the system and framework of supply chain resilience evaluation, and for the supply chain that is not perfect, enterprises must optimize the supply chain if they want to achieve sustainable development. Literature [26] examined supply chain optimization models under risk and uncertainty based on literature review, and the systematic mapping using co-citation and PageRank algorithms presented seven key research themes, and the micro-analysis of these themes provided prospective research directions. Literature [27] provides recent research on natural gas supply chain resilience, the difference between supply and demand caused by disruptions in the supply chain network, reveals the shortcomings of the current research and shows that optimization is an effective way to achieve resilience consistently in supply chain production, storage and transportation activities. Literature [28] introduces a multi-portfolio approach and a scenario-based stochastic mixed-integer planning model for optimizing supply chain operations under cascading effects, revealing that resilience measures used to mitigate the impact of disruptions in a specific region can be effectively applied to mitigate the impact of shocks to which they are subjected. Literature [29] examines various theoretical approaches to AI in the context of supply chain optimization, describes ways to achieve higher performance and adaptability, and provides recommendations for the current state of the art and future directions of AI-driven supply chain optimization. Literature [30] developed a two-criteria mixed-integer linear programming model to design a supply chain aiming to maximize profit and supply density, highlighting that the model

performs better compared to traditional profit maximization methods and will assist firms in evaluating the trade-offs between mitigation benefits and costs. Literature [31] describes the main supply chain disruption risks and presents a hybrid framework for exogenous risk management of non-specific supply chains using agent-based modeling, early warning systems, and many objective optimization and option-aware analyses. Literature [32] proposes a resilient supply chain management structure and utilizes artificial intelligence techniques, particularly long and short-term memory networks and particle swarm optimization, to facilitate the field in achieving resilient and efficient supply chain operations that help to improve performance and reduce costs in complex supply chain environments. Literature [33] examined the problem of resilient supply chain recovery by constructing a generalized model based on bounded optimal control theory to determine the type, scope, and timing of response measures, which was shown to be able to help managers cope with risks based on comparing the available recovery options.

The study firstly designs a supply chain resilience evaluation index system based on the theories of resource-based theory and dynamic capability theory, combined with the factor analysis method, and conducts normality and significance tests on the selected indexes. Then, single machine learning model and Stacking principle are analyzed, and for the problems of too direct way of training data generation of meta-model and difficult selection of meta-classifiers, a new Stacking integration algorithm is proposed, which can further improve the performance of the Stacking integration model, and it is applied to supply chain resilience evaluation prediction. Finally, the evaluation and prediction effects of a single machine learning model and this paper's model on the supply chain toughness of the sample enterprises are compared, and relevant suggestions for supply chain toughness optimization are proposed on the results.

2. Selection and screening of supply chain resilience indicators

2.1. Selection of sample indicators

In this paper, the relevant staff of the supply chain department, purchasing department, sales department, warehousing department and organization department of the car company who understand the supply chain of the enterprise were selected to fill in the questionnaire, and at the same time, the relevant practitioners who have a certain degree of understanding of the Z car company were invited to complete the questionnaire, which was mainly distributed to them through the WeChat platform. A total of 103 questionnaires were issued, 70 valid questionnaires, the predictive degree questionnaire will be used to measure the Likert 5-level scale, the data collected will be used to validate the analysis of the evaluation index system, while looking for problems with the questionnaire during the collection process, with the intention of perfecting the construction of the index system and the questionnaire design through the predictive degree of this questionnaire. The initial questionnaire is shown in Table 1.

Table 1. Initial code

Dimension	Code	Index
Predictive power	YZ1	Strategic management
	YZ2	Process operation ability
	YZ3	Risk perception
	YZ4	Information management ability
Responsiveness	XY1	Contingency planning capability
	XY2	Employee response ability
	XY3	Supplier coordination
	XY4	Information interaction ability
Adaptive ability	SY1	Key resource reserves
	SY2	Internal coordination
	SY3	The ability of mutual resources for enterprises
	SY4	Security ability
Recovery capacity	HF1	Supply chain flexibility
	HF2	Corporate financial capacity
	HF3	Social capital level
Learning ability	XX1	Organizational education capability
	XX2	Experience sharing ability
	XX3	Knowledge acquisition ability

2.2. Construction of the evaluation index system

1) Scale reliability analysis. In conducting the reliability analysis, since the number of items under each secondary index is not the same, in order to ensure a reasonable assessment of the data, the scores of the items under each secondary index are averaged and then analyzed for reliability. Second, this paper will use the overall correlation coefficient (CITC) to determine the specific measurement items. Usually, the CITC index is bounded by 0.5, and when the value of this calculated coefficient is less than 0.5, the corresponding index should be considered for deletion. Meanwhile, the Cronbach's alpha coefficient was chosen in this paper as another index to test the reliability of the measurement items, which has become a well-recognized measure of internal consistency and has been widely used to test the reliability of scales due to its intuitive display of the internal consistency of multidimensional scales. In this paper, the CITC index as well as the Cronbach's alpha coefficient were calculated for each dimension, and the results are shown in Table 2.

From the table, it can be seen that the reliability of each question item of the scale is greater than 0.5, which meets the general standard of reliability test, so there is no need to exclude the indicators in the scale. In addition, the alpha value of the predictive ability indicator dimension reached 0.933, and the alpha value of the applicability ability indicator dimension reached 0.933, which is significantly higher than 0.850, indicating that the reliability of this indicator is good; the alpha indicators of the responsiveness, recovery ability, and learning and growth ability dimensions were 0.871, 0.896, and 0.893, respectively, indicating that the indicators have high reliability. In addition, each of the deleted Cronbach's alpha is greater than 0.7, indicating that the reliability of each indicator meets the requirements.

Table 2. CITC index and Cronbach α coefficient

Dimension	Cronbach reliability analysis			
	Code	CITC	Item deleted Cronbach' α	Cronbach' α
Predictive power	YZ1	0.833	0.911	0.933
	YZ2	0.843	0.907	
	YZ3	0.849	0.908	
	YZ4	0.825	0.911	
Responsiveness	XY1	0.775	0.816	0.871
	XY2	0.791	0.811	
	XY3	0.825	0.795	
	XY4	0.552	0.916	
Adaptive ability	SY1	0.826	0.917	0.933
	SY2	0.895	0.891	
	SY3	0.807	0.924	
	SY4	0.833	0.914	
Recovery capacity	HF1	0.865	0.801	0.896
	HF2	0.794	0.863	
	HF3	0.746	0.902	
Learning ability	XX1	0.781	0.835	0.893
	XX2	0.824	0.811	
	XX3	0.756	0.857	

Also, this study examined the overall reliability of the index system, as shown in Table 3. As can be seen from the table, the overall Cronbach's α is 0.914, which indicates that the quality of the reliability of this predictive measure questionnaire is high.

Table 3. Reliability statistics

Cronbach's α	Term number
0.914	18

2) Scale validity analysis. First, the determination of whether the data samples are capable of exploratory factor analysis is carried out. In this paper, the sampling suitability measure (KMO) and Bartlett's ball (Bartlett's) shape test are used for testing, as shown in Table 4. When the KMO value is greater than 0.7, it can be determined that the data samples can be used for factor analysis, and at the same time, the closer the KMO value is to 1, the stronger the correlation type between the variables. Bartlett's spherical test can be used to test the independence between the variables, when the P value is less than 0.05, it means that the variables are not independent of each other, and the common factor can be extracted from them, that is to say, it is judged that the test variables can be used for exploratory factor analysis. As can be seen from the table, the KMO value is 0.872, which is greater than 0.8, and the P value is $0 < 0.5$, which can be judged as the original variable data correlation, and the predictive degree questionnaire in this paper has good validity, and can be carried out in the next step of exploratory factor analysis.

Table 4. KMO and bartlett test

KMO sampling availability number		0.872
Bartlett sphericity test	Approximate card	15782.763
	Freedom	158
	Significance	0.00

3) Exploratory factor analysis. In this paper, we will use principal component analysis to do exploratory factor analysis on each secondary index in the initial index system established above, and the factors with eigenvalues greater than 1 have been extracted and inductively analyzed as shown in Table 5. Finally, five principal factors are extracted as the five dimensions of this paper, and the cumulative explained variance reaches 77.877%, indicating that the five principal factors can effectively explain the information contained in the original 18 variables.

Table 5. Total variance interpretation

Constituent	Initial eigenvalue			Extracting the load of the load			Rotational load squared		
	Total	Percentage of variance	Cumulation %	Total	Percentage of variance	Cumulation %	Total	Percentage of variance	Cumulation %
1	6.348	35.267	35.267	6.348	35.267	35.267	5.456	30.311	30.311
2	3.466	19.256	54.523	3.466	19.256	54.523	4.404	24.467	54.778
3	1.879	10.439	64.962	1.879	10.439	64.962	1.986	11.033	65.811
4	1.295	7.194	72.156	1.295	7.194	72.156	1.635	9.083	74.894
5	1.03	5.722	77.878	1.03	5.722	77.878	0.537	2.983	77.877
6	0.603	3.35	81.228						
7	0.577	3.21	84.438						
8	0.488	2.711	87.149						
9	0.358	1.989	89.138						
10	0.33	1.833	90.971						
11	0.25	1.389	92.36						
12	0.24	1.333	93.693						
13	0.232	1.289	94.982						
14	0.301	1.672	96.654						
15	0.185	1.028	97.682						
16	0.161	0.894	98.576						
17	0.138	0.767	99.343						
18	0.119	0.661	100						

Second, according to the factor fragmentation diagram, when the fold line changes sharply from a steep trend to a smooth trend, i.e., the number of factors corresponding to the number of factors at the intersection of steep and smooth minus 1 is the reasonable range of factor extraction, as shown in Figure 1. In the figure, starting from the 6th factor, the characteristic roots of the factors are all less than 1, which indicates that starting from the 6th factor means that the cause variables cannot be well explained and can be omitted, and it also indicates that 5 factors should be extracted for analysis in this study.

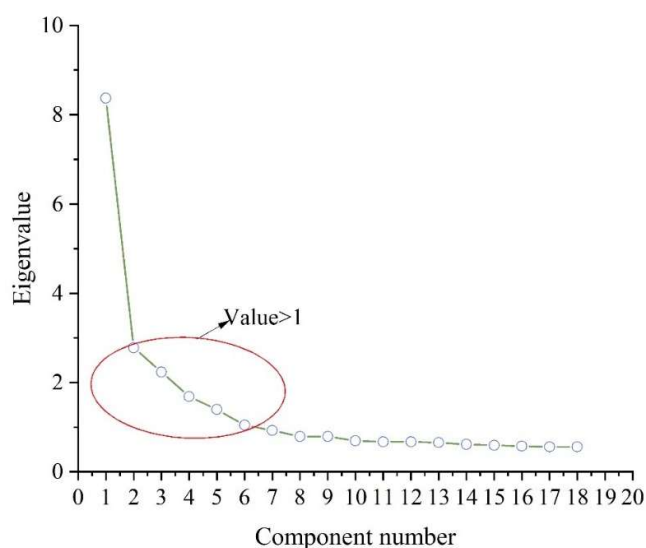


Fig. 1. Factor lithogram

Finally, the maximum variance method was used to rotate the factors orthogonally to extract the high loading factors on each variable, and finally convergence was reached after six iterations to obtain the final rotated component matrix, as shown in Table 6. Among them, the factor loading levels of the four secondary indicators YZ1, YZ2, YZ3 and YZ4 are higher in the first factor range, so the factor one is named as “predictive ability”. The factor loading levels of the four secondary indicators XY1, XY2, XY3, and XY4 were high within the third factor, which was named “responsiveness”. Four secondary indicators, SY1, SY2, SY3 and SY4, had higher factor loadings on the second factor, so factor 2 was named “adaptive capacity”. Three secondary indicators, HF1, HF2 and HF3, had high factor loadings in the fifth factor, so factor V was named “resilience”. The three secondary indicators XX1, XX2, and XX3 have higher factor loadings in the fourth factor, so factor 4 is named “ability to learn and grow”.

Table 6. Rotational composition matrix

Measuring item	Constituent				
	1	2	3	4	5
YZ2	0.835				
YZ3	0.831				
YZ4	0.828				
YZ1	0.812				
SY3		0.891			
SY2		0.838			
SY1		0.775			
SY4		0.762			
XY1			0.904		
XY4			0.897		
XY2			0.895		
XY3			0.775		
XX2				0.845	
XX1				0.839	
XX3				0.832	
HF1					0.868
HF2					0.817
HF3					0.765

4) Determination of evaluation indicators. After the pre-test data reliability and validity analysis, it was realized to test the feasibility of the developed questionnaire, and it was concluded that the initial set dimensions of the indicator system were consistent with the results of the exploratory factor analysis. The finalized evaluation index system of supply chain resilience of Z-vehicle enterprises containing 5 dimensions and 18 indicators was determined, as shown in Table 7.

Table 7. The z-car enterprise supply chain toughness evaluation index system

Dimension	Index
Predictive power	Strategic management capability x1
	Process operations x2
	Risk perception x3
	Information management capability x4
Responsiveness	Emergency planning capability x5
	Employee response ability x6
	Supplier coordination ability x7
	Information interaction ability x8
Adaptive ability	Key resource reserve capability x9
	Internal coordination ability
	The ability of mutual resources to help the enterprise x11
	Security capability x12
Recovery capacity	Supply chain flexibility x13
	Company financial capability x14
	Social capital level x15
Learning ability	Organization education capability x16
	Experience sharing ability x17
	Knowledge acquisition ability x18

2.3. Indicator testing

2.3.1. Normality test. In order to test whether there is a significant difference between the indicators selected in this paper between the ST group and the non-ST group, the indicators need to be tested for significance. According to whether the sample distribution is known or not, the significance test can be divided into parametric test and non-parametric test. Commonly used parametric tests include the t-test, which is applicable to indicators that conform to normal distribution; commonly used non-parametric tests include the Wilcoxon rank sum test and the Mann-Whitney test. Therefore, before conducting the significance test, it is necessary to test the indicators for normal distribution. In this paper, the Shapiro-Wilk test is used to determine whether the indicators meet the normal distribution.

The basic idea of Shapiro-Wilk test is to determine whether the sample data conform to normal distribution by calculating the W value of the sample data and comparing the W value with the W value of normal distribution. The following are the steps of Shapiro-Wilk test:

- 1) Arrange the raw data in ascending order to obtain the order statistic $X_{(1)} \leq X_{(2)} \leq \dots \leq X_{(n)}$.
- 2) Construct the statistic:

$$W = \frac{\left(\sum_{i=1}^n a_i x_i \right)^2}{\sum_{i=1}^n (x_i - \bar{x})^2} \quad (1)$$

where x_i is the i nd order statistic, a constant, which is obtained by looking up the table.

3) Calculate the value of W and compare it with the critical value of W_α .

The original hypothesis of the Shapiro-Wilk test is that the sample data conforms to a normal distribution; the alternative hypothesis is that the sample data does not conform to a normal distribution. By comparing the P-value of the W-value with the significance level, the original hypothesis is accepted when the P-value of the W-value is greater than the significance level α . If the P-value of the W-value is less than the significance level α , the original hypothesis is rejected. Generally, α is taken as 0.05. The results of Shapiro-Wilk test are shown in Table 8 by taking the example of T-3 years of Z car companies. From the results, it can be seen that the P-value of all indicators in T-3 years is less than 0.05, and the same test is performed on the data of T-3, T-4 years, and the P-value is also less than 0.05, which indicates that all indicators in the sample do not obey the normal distribution.

Table 8. T-3 years Shapiro-Wilk test results

Index	Statistic	P value
X1	0.4311	0.000*
X2	0.3822	0.000*
X3	0.4481	0.000*
X4	0.2771	0.000*
X5	0.4228	0.000*
X6	0.6714	0.000*
X7	0.5392	0.000*
X8	0.9734	0.000*
X9	0.2719	0.000*
X10	0.1342	0.000*
X11	0.2018	0.000*
X12	0.0947	0.000*
X13	0.4075	0.000*
X14	0.2436	0.000*
X15	0.1974	0.000*
X16	0.8153	0.000*
X17	0.1375	0.000*
X18	0.2168	0.000*

Note: * indicates significant at the 5% level.

2.3.2. Significance Test. From the above, it can be seen that all indicators did not pass the normality test. Therefore, each indicator is tested for significance using nonparametric tests, and the Mann-Whitney test of nonparametric tests is used in this paper. The test is based on the original hypothesis that there is no significant difference between the two groups of data, and the principle is as follows:

1) The two sets of data are mixed and ranked in order of size. The smallest data rank is 1, the second smallest data rank is 2, and so on, if there is a situation where the data are equal, then take the average of these data sorted as its rank;

2) Find the rank and W_1, W_2 for each of the two samples;

3) Calculate the Mann-Whitney test statistic based on the numbers n_1, n_2 and W_1, W_2 of the two samples:

$$U_1 = n_1 n_2 + \frac{n_1(n_1 + 1)}{2} - W_1 \quad (2)$$

$$U_2 = n_1 n_2 + \frac{n_2(n_2 + 1)}{2} - W_2 \quad (3)$$

Set the smallest of U_1 and U_2 as U and compare it with the critical value of U_α . When $U < U_\alpha$, the original hypothesis is rejected and the alternative hypothesis is accepted, that is, there is a significant difference between the two groups of data. Mann-Whitney test was performed on all indicators and the test results are shown in Table 9. It can also be seen through the results that most of the indicators in the indicators of anticipation, responsiveness, adaptability, resilience, and learning and growth indicators passed the significance test in the years T-2, T-3, and T-4. It indicates that these indicators may have generated differences at the early stage of changes in the resilience of the enterprise supply chain, and belong to the indicators that are more sensitive to changes in the business situation of the enterprise, and the differences persist with the deterioration of the business situation of the enterprise.

Table 9. The Mann-Whitney test results

Index	T-2		T-3		T-4	
	Statistic	P value	Statistic	P value	Statistic	P value
X1	17541	0.000*	17877	0.000*	18224	0.000*
X2	17275	0.000*	17276	0.000*	17351	0.000*
X3	16031	0.000*	17311	0.000*	18262	0.000*
X4	16892	0.000*	18182	0.000*	17973	0.000*
X5	15022	0.000*	18027	0.000*	17674	0.000*
X6	12673	0.000*	15855	0.000*	15853	0.000*
X7	16881	0.000*	17315	0.000*	19247	0.000*
X8	8697	0.000*	18011	0.000*	17832	0.000*
X9	11891	0.000*	18022	0.000*	17832	0.000*
X10	15757	0.000*	23451	0.000*	20753	0.000*
X11	15114	0.000*	20973	0.325	19542	0.000*
X12	16972	0.000*	20816	0.008*	19557	0.000*
X13	21547	0.000*	18754	0.001*	22642	0.145
X14	18865	0.000*	23879	0.465	22176	0.077
X15	19642	0.035*	16975	0.000*	15774	0.000*
X16	23486	0.000*	17137	0.001*	16472	0.000*
X17	8379	0.386	17585	0.000*	19781	0.000*
X18	7467	0.000*	20274	0.000*	19781	0.000*

Note: * indicates significant at the 5% level.

3. Supply chain resilience evaluation model based on integrated learning algorithm

3.1. Single Machine Learning Theory

3.1.1. Logistic regression. Logistic regression is a linear classifier evolved from linear regression and is widely used in classification problems as a common model in financial alerts. The basic form of the model is:

$$P_i = \frac{\exp(\beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_p X_p)}{1 + \exp(\beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_p X_p)} \quad (4)$$

The value of P indicates the probability of SMEs being in financial distress, and in this paper, we set that when $P > 0.5$ SMEs will be in financial distress in the second year; conversely, they will operate normally in the second year.

3.1.2. Support vector machines. A support vector machine is a binary classification model with linear classifiers that are maximally spaced on the feature space [34]. Support vector machines are proposed to effectively solve the problem of linear inability to classify with the following mathematical description:

Given a training set $T = \{(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)\} \in (R^m * \gamma)^n$, where $x_i \in R^m, i = 1, 2, \dots, n$ is a sample and $y_i \in \gamma = \{+1, -1\}, i = 1, 2, \dots, y_i$ is the class label corresponding to sample x_i . The binary classification problem is to find a classification hyperplane $g(x) = \langle w, x \rangle + b, w \in R^m, b \in R$ in R^m such that for any sample x , if $g(x) \geq 0$, then x is judged to belong to class +1, and if $g(x) < 0$, then x is judged to belong to class -1.

Let the set of training examples $[x_i, y_i]$ be composed of input variables $x_i \in R^n$ and categorized values $y_i \in -1, 1, i = 1, \dots, I$, and the support vector of the optimal hyperplane discrete binary decision class rule is determined by the following equation:

$$Y = \text{sign} \left(\sum_{i=1}^N y_i a_i(x, x_i) + c \right) \quad (5)$$

where, Y is the final result, y_i is the set of instances x_i is the vector of indicator attributes representing the classified values of the classification training:

The input variable $x_i, i = 1, \dots, M$ corresponding to each indicator attribute vector is used as the support vector; c and a_i are hyperplane parameters. For a nonlinear discrete instance set, the above equation is transformed into a high-dimensional form as follows:

$$Y = \text{sign} \left(\sum_{i=1}^N y_i a_i k(x, x_i) + c \right) \quad (6)$$

Support vector machines need to introduce kernel functions when dealing with nonlinear data, with the aim of transforming the data dimensions to improve the accuracy of predictions.

The mathematical expression of the kernel function is shown below:

Call $k: R^m \rightarrow R$ a kernel function if there exists a mapping from R^m to the Hilbert space η :

$$\varphi: R^m \rightarrow \eta \quad (7)$$

$$x \rightarrow \varphi(x) \quad (8)$$

Such that $\forall x, z \in R^m$, satisfies $k(x, z) = \langle \varphi(x), \varphi(z) \rangle$, where $\langle \cdot \rangle$ denotes the inner product in space η . In the binary classification problem, φ is regarded as a mapping of features and the image space η of φ is called the feature space. According to a large number of literature studies, it is known that the commonly used kernel functions include four kinds of kernel functions, namely: linear kernel function, polynomial kernel function, radial basis kernel function, and Sigmoid kernel function.

3.1.3. Decision trees. The decision tree algorithm is capable of generating a tree-shaped classification model based on an unordered training set of raw samples that can perform both classification and regression tasks [35]. Whichever feature the model utilizes to make a category judgment is recorded in the leaf nodes of the decision tree, and the information gain, gain ratio, and Gini coefficient are divided according to the order of the 3 node emissions.

Information gain = sample entropy - sum of entropy of all tested attributes.

Where entropy is the sample set purity, theoretically the smaller the entropy value the higher the purity of the sample set, the entropy formula is shown below:

$$\text{Ent}(A) = - \sum_{m=1}^M p_m \log_2 p_m \quad (9)$$

In the above equation, p_m is the percentage of samples in category m .

Information entropy refers to the gain effect of the attributes brought to the test on the purity of the sample set, i.e., the enhancement effect on the purity of the samples. Contrary to the entropy value, the larger the information entropy is, the better it is. The information entropy formula is shown below:

$$Gain(A, a) = Ent(A) - \sum_{w=1}^W \frac{|A^w|}{|A|} Ent(A^w) \quad (10)$$

In the above equation, A^w is the set of samples that satisfy a certain test attribute.

Common decision tree algorithms include CLS algorithm, ID3 algorithm, C4.5 algorithm, SLIQ algorithm, and CART algorithm. The CART algorithm, for example, utilizes the structure of a spanning binary tree to perform the task of binary recursive partitioning, with the advantages of fast computing speed, high computing efficiency, and no restriction on the probability distributions of the predictor variables and the objective function is the algorithm that is used with high frequency in the decision tree. The CART algorithm determines the effect of attribute partitioning based on the magnitude of the Gini coefficient (Gini), and when the smaller the Gini represents the better the classification effect.

$$Gini(T) = 1 - \sum P_J^2 \quad (11)$$

In the above Gini formula, P_J is the probability that category J occurs in sample T .

3.1.4. Plain Bayesian. The naïve Bayesian classification algorithm is based on the Bayesian principle, which mainly explains the problem of "inverse probability", that is, the classification is speculated on the basis of not knowing the whole picture of the data, and then the results are modified through subsequent experiments, which is more in line with the actual situation than the general statistical method of "forward probability" [36]. The plain Bayesian algorithm differs from other machine learning algorithms in that the generative method is utilized to find out the joint distribution probability $P(X, Y)$ of the target classification Y and feature X after which the result is derived from Equation $P(Y | X) = P(X, Y) / P(X)$.

The simple Bayesian conditional probability formula is:

$$P(a_i | x) = \frac{P(a_i)P(x | a_i)}{P(x)} \quad (12)$$

In the above equation $x(x_1, x_2, \dots, x_m)$ is the feature vector, $a(a_1, a_2, \dots, a_n)$ is the category vector, $P(a_i | x)$ is the conditional probability of each category vector using the Bayesian formula to estimate the feature vector of the training set, only when $P(a_i | x)$ takes the maximum value of the data x belongs to the a , $P(x)$ denotes the evidence factor used for normalization, so the plain Bayesian algorithm only needs to estimate the $P(a_i)$ and $P(x | a_i)$ where $P(x | a_i)$ is the likelihood probability and $P(a_i)$ is the prior probability.

3.2. Integrated Learning Approach

3.2.1. Principles of Stacking Integration Algorithm. Stacking integration algorithm is a learning-based integration strategy that learns the same dataset through different base modeling algorithms, and then fuses the learned results with a metamodel, as shown in Fig. 2. The whole framework is mainly divided into two parts, the first part is processed through Tier-1 and the second part is processed through Tier-2. Before Tier-1 processing generally bootstrap sampling to divide each training set, that is, k-fold division of the training set. The base models in each Tier-1 are sequentially trained and decided on the divided training set, each time training T-1 copies of them, and predicting the remaining 1. All base classifiers in the Tier-1 set are processed according to this method to obtain the new data set needed for Tier-2.

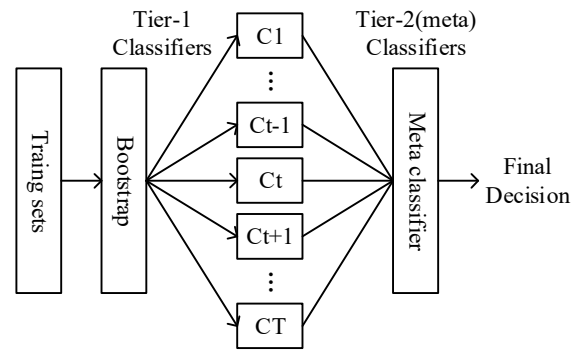


Fig. 2. The overall framework of the stacking

Summarizing the above method theory, then a basic Stacking framework will perform the following operations:

- 1) Choose the base model. We can have 4 basic algorithmic models: logistic regression, support vector machine, decision tree, and plain Bayes.
- 2) Divide the training set into non-crossing k copies, $k=5$ in the text. we label as train1 to train5.
- 3) Starting from train1 as the validation set, model using train2 to train5 as the training set, then predict train1 and keep the results; then, using train2 as the validation set, train the model using train1, train3 to train5, predict train2 and keep the results; and so on until we predict train1 to train5 is predicted once each;
- 4) During the five models built above, each model predicts the Testing dataset separately and ultimately retains the results of these five columns, which are then averaged as a stack transformation of the first base model on the Testing data.
- 5) Select the second base model and repeat the above operations 2-5 to again get one Stacking transformation of the entire dataset of Training in the second base model.
- 6) And so on. With several base models, several new columns of feature representations are generated for the entire Training dataset. Similarly, there will be several new columns of feature expressions for Testing.
- 7) LR is generally used as the second layer of the model for modeling predictions.

3.2.2. Improvement strategies. Considering the learning idea in Stacking, and the fact that the attribute value of each new attribute in traditional Stacking is the predicted value of the target Y , we consider that there is a strong correlation between the feature attributes X , and between X and the target value Y . Multi-classifier fusion can improve the accuracy (bias) and stability (variance) of the model, according to the simple and effective law of "Occam's Razor", we choose four algorithms, namely, logistic regression, plain Bayes, support vector machine and decision tree, as the initial base classifiers in the meta-classifiers, to ensure that the model predictive ability while training the model as fast as possible. The fusion strategy uses a weighted approach to generate a multi-classifier fusion model.

Logistic regression, Improved Gaussian Simple Bayes and Support Vector Machine are used to predict the probability that instance x is category 1 and category 0 respectively. In order to prevent "abnormally" large or small probabilities from dominating the average result, the probabilities are smoothed by multiplying them by a weight, and then averaged to obtain the final probability value and output the corresponding probability value. Finally, the smoothed probability is averaged to get the final probability value and output the corresponding category. The process of weight calculation is as follows:

- 1) First calculate the number of misclassified samples LR_{n_k} , GNB_{n_k} and SVM_{n_k} for each model on the validation set for the 5-fold validation.
- 2) Calculate the corresponding weight values for each model according to Equation (13):

$$weight_model = \frac{avg(LR_{n_k}) + avg(GNB_{n_k}) + avg(SVM_{n_k})}{avg(model_{n_k})} \quad (13)$$

4. Supply chain resilience evaluation and optimization

4.1. Indicators for model evaluation

Model evaluation metrics are numerical metrics used to measure model performance and prediction accuracy, which are usually used to assess the effectiveness of machine learning models and compare the performance of different models. In this paper, mean absolute error (MAE), root mean square error (RMSE), and coefficient of determination (R^2) are chosen as the evaluation indexes.

1) The mean absolute error is the average of the absolute value of the difference between the predicted value and the true value, reflecting the size of the average error of the model, and the smaller the MAE is, the better the predictive ability of the model, which is calculated by the formula:

$$MAE = \frac{1}{n} \sum_{i=1}^n |y_i - \hat{y}_i| \quad (14)$$

2) Root Mean Square Error (RMSE) is obtained by averaging and taking the square root of the squared difference between the predicted value and the true value. The smaller the RMSE is, the more accurate the results predicted by the model, and is calculated by the formula:

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2} \quad (15)$$

3) The coefficient of determination (R^2) is used to assess the degree of model fit by calculating the proportion of variance in the data explained by the model by comparing the difference between the model's predicted and actual values. The value of the coefficient of determination ranges from 0 to 1. The closer the value is to 1, the better the model fits the data, while the closer the value is to 0, the worse the model fits the data, and the formula is:

$$R^2 = 1 - \frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{\sum_{i=1}^n (y_i - \bar{y})^2} \quad (16)$$

4.2. Single forecasting model

4.2.1. Logistic regression evaluation. After establishing and training the best parameter Logistic regression model and applying the test set, the evaluation indexes are obtained as shown in Table 10 and the predictive fit graph as shown in Figure 3, where green represents the real supply chain toughness and blue represents the supply chain toughness predicted by the model. From the coefficient of determination of 45.6785% and combined with the predictive fit graph, it can be seen that the logistic regression model fits the effect generally.

Table 10. Logistic regression model evaluation results

Evaluation index	MAE	RMSE	R^2
results	83.6539	145.284	0.456785

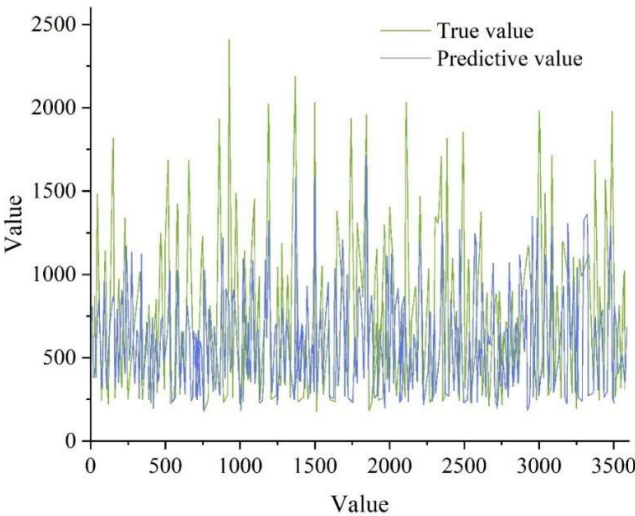


Fig. 3. The regression model is fitted with the rendering

4.2.2. Support vector machine evaluation. After establishing and training the best parameterized support vector regression model and applying it to the test set, the evaluation indexes are obtained as shown in Table 11 and the prediction fitting graph as shown in Figure 4, where green represents the real supply chain toughness and blue represents the supply chain toughness predicted by the model. It can be seen that the evaluation indicators are greatly improved compared with the results of the logistic regression model, and the goodness of fit reaches 92.2254%.

Table 11. Support vector model evaluation index results

Evaluation index	MAE	RMSE	R^2
results	47.3251	85.3822	0.922254

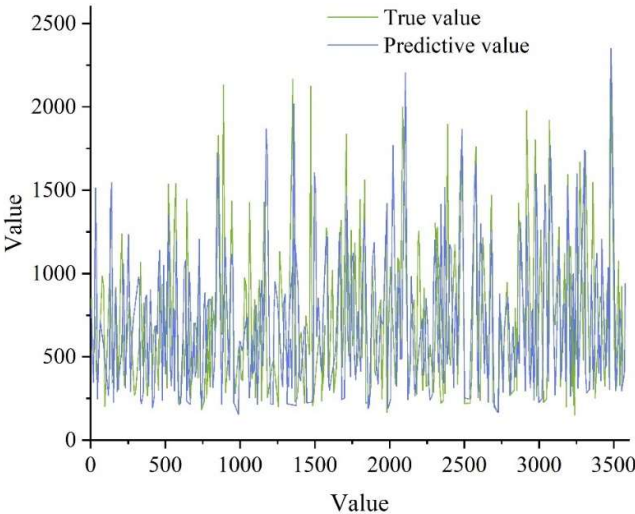


Fig. 4. Support vector machine fitting rendering

4.2.3. Decision tree evaluation. After establishing and training the best parameter decision tree model and applying it to the test set, we get the evaluation indexes as shown in Table 12 and the prediction fitting diagram as shown in Fig. 5, with green representing the real supply chain toughness and blue representing the supply chain toughness predicted by the model. It can be seen that the evaluation indexes compared with the Logistic regression model and the support vector machine model results have been greatly improved, and the goodness of fit reaches 93.5526%.

Table 12. Decision tree model evaluation results

Evaluation index	MAE	RMSE	R^2
Results	31.6647	60.2281	0.935526

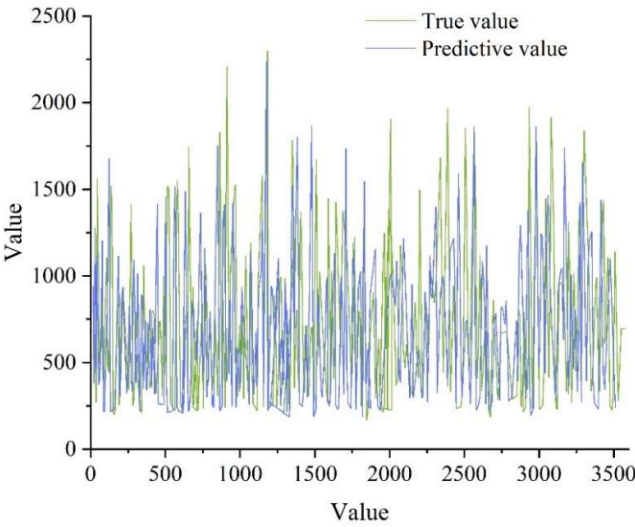


Fig. 5. The decision tree model is fitting the rendering

4.2.4. Plain Bayesian evaluation. After building and training the best parametric Bayesian model and applying it to the test set, the evaluation metrics are obtained as shown in Table 13 and the predictive fit graph as shown in Figure 6, with green representing the real supply chain toughness and blue representing the model-predicted supply chain toughness. Comparing the first three models, the fit of the plain Bayesian model reaches 94.4783%, which is a better fit and is the optimal model.

Table 13. The beeses model is evaluated

Evaluation index	MAE	RMSE	R^2
Results	27.5849	53.0851	0.944783

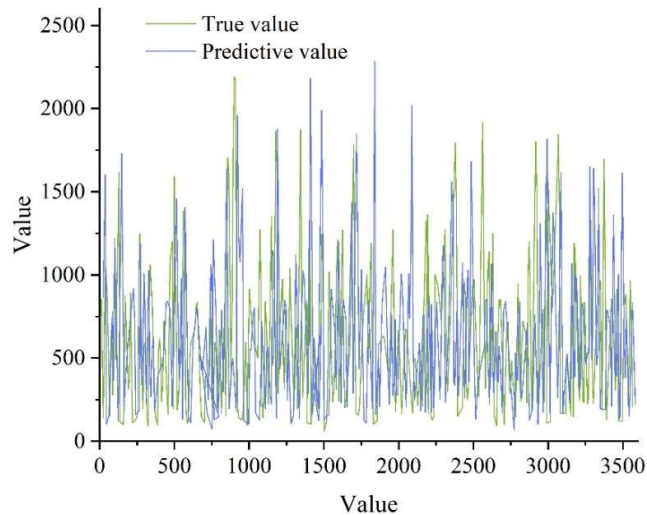


Fig. 6. The beeses model is fitted with a rendering

4.3. Improving the Stacking Integration Model

4.3.1. Primary Learner Selection. Through the construction of four models in the previous section, the prediction results of seven models are summarized, and the results are shown in Table 14. According to the fitting effect graph of a single model and combined with the results summarized by the model evaluation index, it can be seen that the logistic regression model and the prediction results have a large gap with the prediction results of the other models, and the prediction results are not good, so the three models of Support Vector Machine, Decision Tree, and Plain Bayes are selected as the primary learners of the Stacking integrated model.

Table 14. The single model evaluation index results are summarized

Model	MAE	RMSE	R^2
Logistic regression	83.6539	145.284	0.456785
Support vector machine	47.3251	85.3822	0.922254
Decision tree	31.6647	60.2281	0.935526
Plain beth	27.5849	53.0851	0.944783

4.3.2. Secondary Learner Selection. In this paper, four single models used in section 4.2 are selected, and the coefficient of determination is used as the evaluation index to compare the effect of each model as a secondary learner on the test set, and the best model is selected, and the results are shown in Table 15. The table shows that the results of each model are very good and the difference is not very big, the logistic regression model with the highest coefficient of determination is selected as the secondary learner of Stacking integration model.

Table 15. Different model test sets

	Logistic regression	Support vector machine	Decision tree	Plain beth
R ²	0.943521	0.932412	0.943158	0.958729

4.3.3. **Stacking Integration Model Construction.** In this paper, three models, namely, Support Vector Machine, Decision Tree, and Plain Bayes, are selected as the primary learners of Stacking integrated models, and five-fold cross-validation is used to predict the original training set respectively, and the five model predictions are stacked horizontally and added with the sample labels of the original training set to get the new training set, and then the trained models are predicted to the original test set, and the five model predictions are stacked horizontally and add the original test set sample labels to get a new test set. The results obtained by the secondary learner logistic regression model after training the new training set and predicting the new test set are used as the final prediction results.

4.3.4. **Analysis of model prediction results.** From the results in Table 14, it can be seen that the seven single models have the best results of the plain Bayesian model, and the predictions of this model are compared with the predictions of the Stacking integrated model, and the results are shown in Table 16. The table shows that the Stacking integrated model performs better than the plain Bayes in all the metrics, which makes it better than any of the single models. The Stacking integrated model is slightly better than the plain Bayesian model in terms of the coefficient of determination, while the effect is improved by 2.2307 in terms of the mean absolute error and 1.3268 in terms of the root-mean-square error, indicating that the Stacking integrated model outperforms the plain Bayesian model in terms of supply chain resilience prediction.

Table 16. Comparison of the simple bees model with stacking integration model

model	MAE	RMSE	R ²
Plain beth	27.5849	53.0851	0.944783
Stacking integration model	25.3542	51.7583	0.953157

To further illustrate the accuracy of the Stacking integrated model in second-hand house price prediction, this paper chooses the relative error to test the sample bias of a single test set. The relative error is the ratio of the error between the predicted value and the true value to the true value, the closer the value is to 0, the better the model prediction effect, calculated as follows

$$\delta_i = \frac{|y_i - \hat{y}|}{y_i} \quad (17)$$

The prediction results of the Stacking integration model are organized and obtained as shown in Table 17. Organizing the prediction results, it is found that there are 3134 samples with relative errors within 10%, accounting for 87.58%, 3467 samples with relative errors within 16%, accounting for 94.41%, and 3685 samples with relative errors within 23%, accounting for 98.78%. Therefore, it shows that the Stacking integration model established in this paper has good prediction effect and high accuracy, which has certain value and significance for the research of supply chain toughness prediction, and can provide scientific reference basis for enterprises.

Table 17. The stacking integration model predicts the results

Product number	True value	Predictive value	Relative error
A98327	341	339.0841	0.55288%
A78562	557	547.6374	1.0584%
B84974	845	892.3821	6.62321%

B80562	352	345.3947	3.47362%
.....			
C73386	2001	2067.9162	2.73415%
C75997	986	915.33792	7.4482%
EY6482	574	549.7139	4.87213%
GY3784	256	274.47158	9.55491%

4.4. Supply chain resilience optimization recommendations

1) Ease financing pressure and broaden funding channels. Reducing the financial constraints on innovation investment is the key to the effective implementation of R&D activities by enterprises. From the perspective of enterprises, on the one hand, enterprises should focus on building and strengthening the supply chain resilience, utilize the advantages of supply chain resilience on enterprise financing, broaden the channels of funding sources, and provide a solid guarantee for the enterprises to carry out R&D activities; on the other hand, enterprises can expand their competitive advantages by strengthening the cooperation and exchange with upstream and downstream enterprises in the supply chain, maintain the high-quality growth of performance and lower financial fluctuation risks, and strengthen the construction and maintenance of corporate reputation. construction and maintenance of the company's reputation, and release positive signals to the outside world, thus alleviating the pressure of the financing constraints imposed. From the government's perspective, on the one hand, the government plays a crucial role in promoting the innovative development of high-tech enterprises. In order to eliminate the "financial obstacles" faced by enterprises in carrying out R & D activities, the government should actively play a policy role to broaden the multi-directional financing channels. For example, it should introduce financial and tax incentives for enterprises and provide national subsidies for science and innovation, so as to alleviate the pressure of financing constraints arising from the constraints imposed on enterprises by other entities. On the other hand, the government should provide a fair financing environment. It should build a financing risk prevention and compensation mechanism and encourage enterprises to finance through multiple ways in order to break the limitations of traditional financing channels. In addition to the existing banks and other financing institutions, enterprises can actively explore the use of supply chain financing, commercial credit financing, enterprise group financing and other new financing methods, so as to make the financing methods more diversified and create a better environment for independent innovation.

2) Enhance market share and consolidate market position. Enterprises can improve their market position by continuously strengthening their product share in the market. Compared with enterprises with higher market position, enterprises in inferior position have lower motivation to innovate, therefore, enterprises should actively build core competitiveness and improve their market position. For the leading or advantageous enterprises in the industry, their position is solid and their resources are relatively abundant, so they should actively explore feasible paths for the positive effects of supply chain resilience; for enterprises with weak competitiveness in the industry, they should speed up the establishment of enterprise collaboration and interconnection alliances to jointly deal with market challenges and risks, and to enhance their own risk-resistant ability, so as to promote their own performance growth and improve their own market position. In exploring these paths, managers need to fully understand and adapt to the laws of operation in a market economy. In addition, the market environment and innovation ecology still need to be further improved. The government should create a more open and fair competitive environment, and promote enterprises to realize the virtuous cycle of supply chain resilience, market position and R&D investment intensity, so that in the face of major emergencies, those suppliers and customers who maintain a long-term relationship with the

enterprise can play a more significant financial risk buffer, turn risks into opportunities, and improve their competitiveness.

3) Deepen the reform of state-owned enterprises and promote regional integration. Deepening the reform of state-owned enterprises should accelerate the reform process and finely coordinate the balanced relationship between enterprise operations and supply chain resilience. In non-essential areas, reduce the degree of non-obvious government guarantees for state-owned enterprises, play the role of public product attributes of the supply chain to make up for short boards and forging long boards, and promote enterprise innovation and high-quality development.

5. Conclusion

Analyze the toughness level of industrial chain supply chain from various aspects and angles, and then put forward countermeasure suggestions to accelerate the construction of industrial chain supply chain ecosystem. The study constructs supply chain toughness evaluation index system based on factor analysis, and selects logistic regression, support vector machine, decision tree and plain Bayes as meta-classifiers to evaluate and predict the toughness level of industrial chain supply chain using the improved Stacking integration algorithm. Taking MAE, RMSE and R^2 as evaluation indexes, the MAE of Stacking integrated model is obtained as 25.3542, RMSE as 51.7583 and R^2 as 0.953157, which are all better than the plain Bayesian model with the best prediction effect, which illustrates the accuracy and superiority of Stacking integrated model.

Funding

This work was supported by the Project of Chengdu Key Research Base of Philosophy and Social Sciences (Project No.2022103, 2022101).

References

- [1] Eggers, S. (2021). A novel approach for analyzing the nuclear supply chain cyber-attack surface. *Nuclear Engineering and Technology*, 53(3), 879-887.
- [2] Hildebrandt, M., Lamshöft, K., Dittmann, J., Neubert, T., & Vielhauer, C. (2020, June). Information hiding in industrial control systems: An OPC UA based supply chain attack and its detection. In *Proceedings of the 2020 ACM Workshop on Information Hiding and Multimedia Security* (pp. 115-120).
- [3] Polatidis, N., Pavlidis, M., & Mouratidis, H. (2018). Cyber-attack path discovery in a dynamic supply chain maritime risk management system. *Computer Standards & Interfaces*, 56, 74-82.
- [4] Gunasekaran, A., Subramanian, N., & Rahman, S. (2015). Supply chain resilience: role of complexities and strategies. *International Journal of Production Research*, 53(22), 6809-6819.
- [5] Zouari, D., Ruel, S., & Viale, L. (2021). Does digitalising the supply chain contribute to its resilience?. *International Journal of Physical Distribution & Logistics Management*, 51(2), 149-180.
- [6] Tukamuhabwa, B. R., Stevenson, M., Busby, J., & Zorzini, M. (2015). Supply chain resilience: definition, review and theoretical foundations for further study. *International journal of production research*, 53(18), 5592-5623.
- [7] Wieland, A., & Durach, C. F. (2021). Two perspectives on supply chain resilience. *Journal of Business Logistics*, 42(3), 315-322.
- [8] Scholten, K., & Schilder, S. (2015). The role of collaboration in supply chain resilience. *Supply Chain Management: An International Journal*, 20(4), 471-484.

- [9] Pettit, T. J., Croxton, K. L., & Fiksel, J. (2019). The evolution of resilience in supply chain management: a retrospective on ensuring supply chain resilience. *Journal of business logistics*, 40(1), 56-65.
- [10] Ivanov, D. (2021). *Introduction to supply chain resilience: Management, modelling, technology*. Springer Nature.
- [11] Novak, D. C., Wu, Z., & Dooley, K. J. (2021). Whose resilience matters? Addressing issues of scale in supply chain resilience. *Journal of Business Logistics*, 42(3), 323-335.
- [12] Behzadi, G., O'Sullivan, M. J., & Olsen, T. L. (2020). On metrics for supply chain resilience. *European Journal of Operational Research*, 287(1), 145-158.
- [13] Zhao, N., Hong, J., & Lau, K. H. (2023). Impact of supply chain digitalization on supply chain resilience and performance: A multi-mediation model. *International Journal of Production Economics*, 259, 108817.
- [14] Agarwal, N., Seth, N., & Agarwal, A. (2022). Evaluation of supply chain resilience index: a graph theory based approach. *Benchmarking: An International Journal*, 29(3), 735-766.
- [15] Singh, C. S., Soni, G., & Badhotiya, G. K. (2019). Performance indicators for supply chain resilience: review and conceptual framework. *Journal of Industrial Engineering International*, 15(Suppl 1), 105-117.
- [16] Aguila, J. O., & ElMaraghy, W. (2019). Supply chain resilience and structure: An evaluation framework. *Procedia manufacturing*, 28, 43-50.
- [17] Zhou, Z., & Hosseini, S. (2023). Evaluating supply chain resilience: a resilience capacity perspective. *International Journal of Logistics Systems and Management*, 46(3), 333-355.
- [18] Haralambides, H., & Gujar, G. (2023). The 'new normal', global uncertainty and key challenges in building reliable and resilient supply chains. *Maritime Economics & Logistics*, 25(4), 623-638.
- [19] Sawyerr, E., & Harrison, C. (2020). Developing resilient supply chains: lessons from high-reliability organisations. *Supply Chain Management: An International Journal*, 25(1), 77-100.
- [20] Peters, E., Knight, L., Boersma, K., & Uenk, N. (2023). Organizing for supply chain resilience: a high reliability network perspective. *International Journal of Operations & Production Management*, 43(1), 48-69.
- [21] Kumar, D., Soni, G., Mangla, S. K., Liao, J., Rathore, A. P. S., & Kazancoglu, Y. (2024). Integrating resilience and reliability in semiconductor supply chains during disruptions. *International Journal of Production Economics*, 276, 109376.
- [22] Wu, A., Sun, Y., Zhang, H., Sun, L., Wang, X., & Li, B. (2023). Research on resilience evaluation of coal industrial chain and supply chain based on interval type-2F-PT-TOPSIS. *Processes*, 11(2), 566.
- [23] Wu, A., Li, P., Sun, L., Su, C., & Wang, X. (2024). Evaluation Research on Resilience of Coal-to-Liquids Industrial Chain and Supply Chain. *Systems*, 12(10), 395.
- [24] Patidar, A., Sharma, M., Agrawal, R., & Sangwan, K. S. (2023). Supply chain resilience and its key performance indicators: an evaluation under Industry 4.0 and sustainability perspective. *Management of Environmental Quality: An International Journal*, 34(4), 962-980.
- [25] Chen, X., Xi, Z., & Jing, P. (2017). A unified framework for evaluating supply chain reliability and resilience. *IEEE Transactions on Reliability*, 66(4), 1144-1156.
- [26] Suryawanshi, P., & Dutta, P. (2022). Optimization models for supply chains under risk, uncertainty, and resilience: A state-of-the-art review and future research directions. *Transportation research part e: logistics and transportation review*, 157, 102553.
- [27] Emenike, S. N., & Falcone, G. (2020). A review on energy supply chain resilience through optimization. *Renewable and Sustainable Energy Reviews*, 134, 110088.

- [28] Sawik, T. (2022). Stochastic optimization of supply chain resilience under ripple effect: A COVID-19 pandemic related study. *Omega*, 109, 102596.
- [29] Abaku, E. A., Edunjobi, T. E., & Odimarha, A. C. (2024). Theoretical approaches to AI in supply chain optimization: Pathways to efficiency and resilience. *International Journal of Science and Technology Research Archive*, 6(1), 092-107.
- [30] Kungwalsong, K., Mendoza, A., Kamath, V., Pazhani, S., & Marmolejo-Saucedo, J. A. (2022). An application of interactive fuzzy optimization model for redesigning supply chain for resilience. *Annals of Operations Research*, 315(2), 1803-1839.
- [31] Anagnostou, A., Mintram, K., & Taylor, S. J. (2024, December). Supply Chain Resilience Optimization with Agent-Based Modeling (SCROAM): A Novel Hybrid Framework. In *2024 Winter Simulation Conference (WSC)* (pp. 1209-1220). IEEE.
- [32] Xu, J., & Bo, L. (2024). Optimizing Supply Chain Resilience using Advanced Analytics and Computational Intelligence Techniques. *IEEE Access*.
- [33] Khamseh, A., Teimoury, E., & Shahanaghi, K. (2021). A new dynamic optimisation model for operational supply chain recovery. *International Journal of Production Research*, 59(24), 7441-7456.
- [34] Qiyu Yang, Congzi Xia, Yuanwu Cui, Xiaoqiuyan Zhang, Tianyu Zhang, Yong Huang... & Min Hu. (2025). Fusion of terahertz spectroscopy and Raman spectroscopy combined with support vector machine to distinguish different pericarpium citri reticulatae species. *Heliyon*, 11(6), e43136-e43136.
- [35] Yadong Zhang, Shaoping Wang, Enrico Zio, Chao Zhang, Hongyan Dui & Rentong Chen. (2025). Model-guided system operational reliability assessment based on gradient boosting decision trees and dynamic Bayesian networks. *Reliability Engineering and System Safety*, 259, 110949-110949.
- [36] Yulin He, Guiliang Ou, Philippe Fournier Viger & Joshua Zhexue Huang. (2025). Attribute grouping-based naive Bayesian classifier. *Science China Information Sciences*, 68(3), 132106-132106.