

Using a hybrid deep neural network and TOPSIS model for small, medium and micro enterprises access to credit-based loans across theoretical reputation strategy

Wanting Chen

Xuanyi Wu

ZY Chen

Yahui Meng*

Ruei-Yuan Wang[†]

School of Science

Guangdong University of Petrochemical Technology

Maoming

Guangdong 525000

China

Timothy Chen[§]

Division of Engineering and Applied Science

California Institute of Technology

Pasadena

CA 91125

United States

Abstract

Small, medium and micro enterprises started late in China, with poor transparency of financial information and relatively weak stability of profitability and asset strength, which makes commercial banks need to bear more risks when providing loans to small, medium and micro enterprises than large enterprises. When there is no credit record of some small and medium-sized micro enterprises in commercial banks, it will increase the risk that banks need to bear when they lend to these small and medium-sized micro enterprises without credit record, and it will also increase the difficulty of small and medium-sized micro enterprises' credit loans in commercial banks. It is a big problem that how to accurately predict and evaluate the small and medium-sized micro enterprises with unknown reputation, and then reduce the risk pressure of banks to provide loan services to these small and medium-sized micro enterprises with unknown reputation. The purpose of this project is to predict and evaluate the credit risk of enterprises with unknown reputation, and then evaluate the

* E-mail: mengyahui@gdupt.edu.cn (Corresponding Author)

† E-mail: rueiyuan@gmail.com

§ E-mail: t13929751005@gmail.com

comprehensive strength of small, medium and micro enterprises that need loans according to the TOPSIS comprehensive evaluation model, so as to calculate the interest rate of bank loans to each type of enterprises, and provide a credit strategy for commercial banks when the total amount of bank credit is 100 million yuan. We use DNN intensive neural network algorithm to list the important parameters of enterprise capability, train a large number of raw data covered by these parameters, so as to evaluate the credit rating of enterprises, and then use TOPSIS comprehensive evaluation model to evaluate the strength of enterprises. Finally, we use k-means clustering algorithm to classify these 203 enterprises and make the best credit decision.

Subject Classification: 93E03, 93A14, 93C15, 49J15.

Keywords: AHP, TOPSIS model, DNN algorithm, K-means clustering algorithm.

1. Introduction

Small and medium-sized micro enterprises, as an indispensable part of China's real economy, play an important role in economic growth, improving the employment rate and scientific and technological development, although the asset scale and operation scale of small and medium-sized micro enterprises are relatively small in the market [1]. In the case of capital problems in small and medium-sized micro enterprises, the person in charge of the company generally tends to borrow money from the bank [2-3]. In the aspect of banks, in the process of bank lending, there is a risk that the recovery of loan funds is hindered [4-5]. In order to avoid this kind of risk, banks usually evaluate all kinds of companies according to credit policy, transaction bill information of enterprises and influence of upstream and downstream enterprises [6-7]. Banks will be more inclined to provide loans to enterprises with strong strength and stable supply and demand relationship. Secondly, according to the enterprise's strength, reputation and credit risk assessment to determine whether to lend and loan line, interest rate and term of credit strategy [8].

2. Problem analysis

In fact, due to the large number of small and medium-sized micro enterprises, banks do not have complete credit rating data of all related enterprises. This also requires us to train the existing data, and then predict the credit rating of unknown enterprises, and then give the lowest risk credit strategy on this basis. In the case of total bank credit of 100 million yuan, we give the lowest risk credit strategy for SMEs with unknown credit rating. First of all, the data of known enterprises' credit rating is used as training data. Then, according to the training results, the

data of small and medium-sized micro enterprises without credit rating is preprocessed. Then, the data of small and medium-sized micro enterprises without credit rating is processed and predicted by using DNN dense link neural network algorithm, and then the risk credit rating of 203 enterprises is predicted and evaluated. Finally, we use TOPSIS comprehensive evaluation model to evaluate the comprehensive strength of enterprises, calculate the interest rates of different enterprises, then use k-means clustering algorithm to classify the 203 enterprises, and finally make the credit decision.

3. Model hypothesis

Due to the limitations of TOPSIS comprehensive scoring model, the assumed positive and negative ideal solutions are too ideal, and DNN cannot deal with the changes in time series, the following conditions need to be assumed to meet the calculation requirements of TOPSIS comprehensive scoring model and DNN neural network algorithm

1. Suppose the data given by the title is true and reliable.
2. It is assumed that the enterprise will not close down in the short term.
3. The hypothesis does not consider the influence of internal factors.
4. It is assumed that every economic activity of the enterprise is reasonable and real.
5. Suppose that all enterprises with unknown reputation apply for credit loans to commercial banks at the same time.

4. Data preprocessing

4.1 AHP (*Analytic Hierarchy Process*)

First of all, because we need to normalize the data of enterprises, we need to convert the ABCD of credit rating into a number. Therefore, we use AHP to calculate the most suitable score proportion for ABCD, reduce the comprehensive ranking error after normalization, and make the results more scientific.

4.1.1 building hierarchical structure model

According to their relationship, the goal, the factor (decision criterion) and the object of important decision-making are the middle level and the

lowest level, divided into the highest level, and the hierarchy diagram is drawn. In this project, the decision objective is set as the best score proportion of ABCD, the consideration factors are set as the replacement indexes of a, B, C and D, and the decision objects are set as different replacement indexes of ABCD.

4.1.2 construction of judgment matrix

We should compare factors in pairs rather than comparing all factors together. Meanwhile, the use of relative scale can reduce the difficulty of comparing various important factors with different properties and increase accuracy. For instances, under a certain criterion, we can compare and classify each option based on its importance. The matrix formed by pairwise comparison results is called judgment matrix, and the elements of judgment matrix[9]. The judgment matrix has the following properties

$$b_{ij} = \frac{1}{b_{ji}} \quad (1)$$

$$B = \begin{bmatrix} b_{11} & b_{12} & \dots & b_{1n} \\ b_{21} & b_{22} & \dots & b_{2n} \\ \dots & \dots & \dots & \dots \end{bmatrix} \quad (2)$$

4.1.3 hierarchical single sort and its consistency test

Based on the findings of matrix theory, when the consistency of the judgment matrix cannot be guaranteed, the corresponding eigenvalues of the judgement matrix will also change. As a result, the alteration in eigenvalues of the judgment matrix can be utilized to examine the consistency of the judgments, and the consistency index of the judgment matrix is introduced into AHP. To verify the consistency of human thought during judgment, the consistency index can be recorded as Ci.

$$CI = \frac{\lambda_{\max} - n}{n - 1} \quad (3)$$

When the CI value of a judgment matrix is high, it indicates a significant deviation from complete consistency. In contrast, a low CI value indicates that the judgment matrix is closer to complete consistency.

4.1.4 solution of AHP analytic hierarchy process model

Through the exhaustive method to list a variety of schemes, and through the consistency test, the best scheme layer is shown in Table 2, its

calculation weight vector is: 0.1172, 0.1775, 0.2806, 0.4247, the scheme layer consistency experience $CI = 0.000734439514301262 \approx 0$, has satisfactory consistency. In other words, the best substitution index of credit rating is $a = 100$, $B = 88.28$, $C = 70.53$, $d = 42.47$.

4.2 *problem data processing*

The data in Appendix 2 is preprocessed in two aspects. After the invoice is summarized, the data in Appendix 2 is predicted by DNN neural network algorithm based on the training data of enterprises with existing reputation level in Appendix 1. On the one hand, it preprocesses the data of interest rate, obtains its weight according to AHP, replaces the index value, normalizes the data, calculates the weight of each factor of TOPSIS comprehensive evaluation model according to entropy weight method, and finally establishes TOPSIS comprehensive evaluation model to calculate the corresponding interest rate of different enterprises. On the other hand, we preprocess the data for the calculation of the quota. After summarizing the invoice, we transform it into array, and then reduce the dimension of the data. We use elbow method and contour counting method to find out the best cluster number, and then cluster analysis to find out the corresponding quota of different enterprises.

5. Model establishment and solution

5.1 *Establish TOPSIS Model*

By assuming the negative and positive ideal solutions, the positive and negative ideal solutions and the distance between each sample is calculated, and the relative closeness between the sample and the ideal solutions are obtained, so as to rank the evaluation objects. In other words, the larger the positive ideal solutions are, the higher the cost the bank needs to pay and the lower the bank's profit; the larger the negative ideal solution is, the lower the cost the bank needs to pay and the higher the bank's profit. Then we make a comprehensive score. We calculate the weight according to the credit rating, default, income, supply and demand stability index, void invoice rate, negative invoice rate, positive ideal solution and negative ideal solution, and the weight is calculated according to the entropy weight method [10]. The relevant symbols are shown in Table 1.

Table 1
Symbol description

R	Raw data matrix
m	Sample size
n	Total index
p	Standardization matrix
W	Weight vector
Z	Weighted normalized matrix
D_i^+	The distance of rational thinking
D_i^-	The distance of negative ideal solution
Ci	Closeness
$\sigma(z)$	Activation function
a_1^2, a_2^2, a_3^2	Output value of the second layer
a_1^3	Output value of the third layer

The specific steps are as follows:

The indicators are in the same direction, standardized and weighted. In order to avoid the influence of subjective factors, the standardized matrix P and weight vector w are obtained.

$$Z = (Z_{ij})_{n \times m} = (p_{ij} * w_j) \quad (4)$$

- (2) The weighted normalized matrix Z is obtained. The positive and negative ideal solutions of Z are obtained by multiplying P and W. The whole ideal solutions mean that each index reaches the best points in these samples, and the negative ideal solution means that each index are the worst values in the samples.

$$D_i^+ = \sqrt{\sum_{j=1}^m (Z_{ij} - Z_j^+)^2}, \quad D_i^- = \sqrt{\sum_{j=1}^m (Z_{ij} - Z_j^-)^2}, (i = 1, \dots, n) \quad (5)$$

- (3) Thees distances between samples and negative and positive ideal solutions is calculated.

$$C_i = \frac{D_i^-}{D_i^+ + D_i^-} \quad (6)$$

5.2 Establishment of DNN dense link neural network prediction model

Suppose that the activation function of selection is $\sigma(z)$. The output values of hidden layer and output layer are σ . For the three-layer DNN in the figure below, the output of the next layer is calculated by using the output of the previous layer. For the output of the second layer a_1^2, a_2^2, a_3^2 , we have

$$a_1^2 = \sigma(z_1^2) = \sigma(w_{11}^2 x_1 + w_{12}^2 x_2 + w_{13}^2 x_3 + b_1^2) \quad (7)$$

$$a_2^2 = \sigma(z_2^2) = \sigma(w_{21}^2 x_1 + w_{22}^2 x_2 + w_{23}^2 x_3 + b_2^2) \quad (8)$$

$$a_3^2 = \sigma(z_3^2) = \sigma(w_{31}^2 x_1 + w_{32}^2 x_2 + w_{33}^2 x_3 + b_3^2) \quad (9)$$

For the output of the third layer We have:

$$a_1^3 = \sigma(z_1^3) = \sigma(w_{11}^3 x_1 + w_{12}^3 x_2 + w_{13}^3 x_3 + b_1^3) \quad (10)$$

If there are m neurons in the $L-1$ layer, the output of the j -th neuron in the L layer will be reduced, We have:

$$a_j^l = \sigma(z_j^l) = \sigma(\sum_{k=1}^m w_{jk}^l a_k^{l-1} + b_j^l) \quad (11)$$

5.3 Problem solving

5.3.1 Quantitative analysis of problem data

After preprocessing the data in Appendix 2 and summarizing the invoice statistics, read the credit rating data of question 1, mark it well, and then merge the enterprises in Appendix 1 and Appendix 2. According to the principal component analysis, as shown in Figure 1, we can determine that the best dimensionality reduction is 8. After training all the enterprises in Annex I, we can get the training accuracy and get the reputation prediction a index. Then we can establish the DNN dense link neural network model and train the first 100 enterprises with DNN. As shown in Figure 2, it is the error of the last 23 tests, and the error of the training model is within the allowable range. After the training model is built, DNN will predict the data in Appendix II, as shown in Figure 3. It is to detect the accuracy of the results of DNN prediction model. The detection results are close to 1. If the model can be established, then the reputation prediction index B can be obtained. Through statistical calculation, the credit evaluation of other enterprises is predicted.

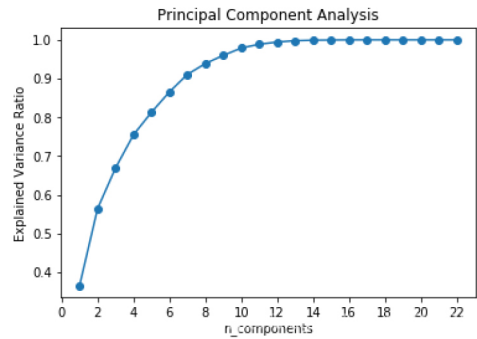


Figure 1
PCA difference preservation drawing

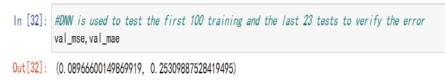


Figure 2
Error analysis of training model

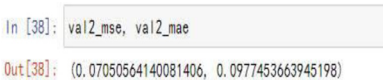


Figure 3
Error analysis of prediction model

	Credit rating	Breach of contract	income	Supply and demand	Void invoice rate	Negative invoice r	Positive ideal s	Negative ideal sol	Comprehensive sort	interest rate
0	0.035137496	0.018998917	0.002750231	0.006143532	0.029372045	0.03027664	0.026115225	0.016231706	0.383303011	204
1	0.035137496	0.020135692	0.005477266	0.006574931	0.029337323	0.03036816	0.025520367	0.016460735	0.392098779	194
2	0.035137496	0.018998976	0.023176421	0.011052348	0.029004334	0.031082648	0.018281283	0.020883668	0.533223396	48
3	0.035137496	0.022759908	0.035137496	0.013235348	0.033879553	0.034877399	0.01546651	0.025765526	0.624890948	2
4	0.023093147	0.017673966	0.015470588	0.024637921	0.033245384	0.034342157	0.014156875	0.022284928	0.611521001	3
5	0.035137496	0.020569256	0.017650796	0	0.031174194	0.032889372	0.025942279	0.018298707	0.413614362	168
6	0.017138157	0.017391365	0.006353992	0.02241941	0.033238771	0.033139242	0.018980457	0.019046426	0.500867392	75
7	0.035137496	0.012222516	0.009747364	0.001624125	0.031610835	0.031216458	0.027199264	0.015819295	0.367731868	231
8	0.035137496	0.014776689	0.011317242	0.005132794	0.033478321	0.034523312	0.02458808	0.016869423	0.406908794	171
9	0.023093147	0.0167710834	0.007899575	0.021772812	0.033188551	0.03448355	0.01802085	0.019677929	0.521977879	53
10	0.035137496	0.014869651	0.010171182	0.001588031	0.032539244	0.031899332	0.02689054	0.016190294	0.375811999	218
11	0.035137496	0.016434116	0.009759207	0.013364994	0.033869654	0.033772404	0.020467227	0.018755239	0.478175937	96
12	0.017138157	0.01897599	0.004504782	0.034501845	0.033851948	0.034688419	0.017519786	0.026204421	0.599311516	12
13	0.023093147	0.018605094	0.00646872	0.011089906	0.032991342	0.03224996	0.022880075	0.014905855	0.394481618	189

Figure 4
result 3 data (part)

After the data of the enterprises whose credit rating has been evaluated is processed twice by TOPSIS model, the brand new data is obtained, as shown in Figure 4. We deal with the results of the data and then there is a detailed analysis. At the same time, according to PCA principal component analysis and PCA difference retention drawing analysis, we can get the result as shown in Figure 8. From figure 5, we can see that when dimension is 6, 91% of the difference can be retained, which is more suitable for the next dimension reduction.

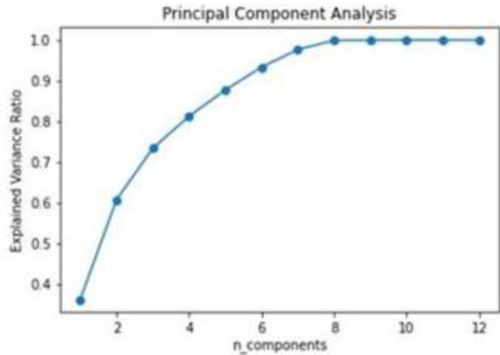


Figure 5

PCA difference retention chart

To reduce the dimension of the preprocessed data, we need to use k-means clustering algorithm to classify the data of result3, and the best number of classification can be obtained by elbow method and contour analysis method. As shown in Figure 6, it is the test of elbow method on the number of classification of all data in result3. As can be seen from Figure 3, when k is from 1 to 10, the decrease of the sum of squares of deviations is quite obvious, while after K is from 10 to 10, the decrease of the sum of squares of deviations is very limited, so the best value of K should be 10 [11]. That is to say, the best number tested by elbow method is 10.

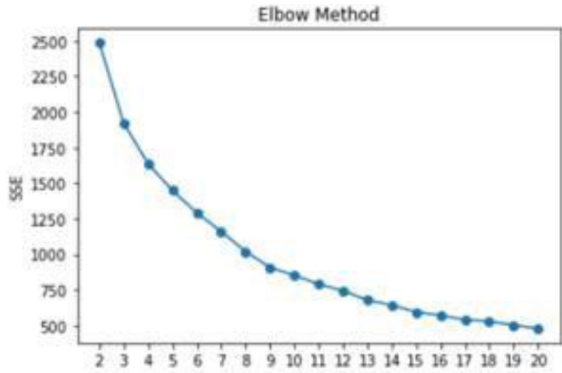


Figure 6

Elbow rule to calculate the best cluster number

As shown in Figure 7, it is the test of classification number of all data in result3 by contour coefficient method. The standard of contour

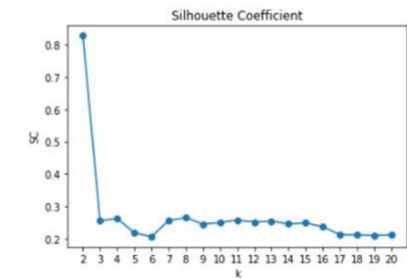


Figure 7

The best clustering number of contour coefficient calculation

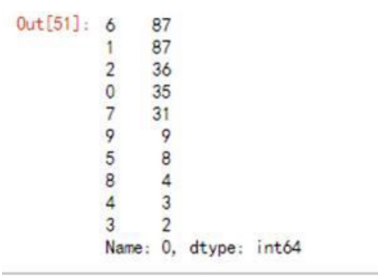


Figure 8

Results of K-means clustering algorithm

	Enterprise	cc	input	inv	Effective amount	Effective input	tax	sales	inv	Effective amount	Effective output	to income	Stability index	Pin stability index	Supply and demand	Void invoice rate	Negative invoice sort	classification	
124	E124	17411	78552553.5	58877291.89	1293	741780367.8	66657380.34	-35765049.06	0.127702775	0.323279196	0.24580524	0.125100303	0.020476803	167	3				
125	E125	20298	93238334.8	68816562.88	1556	941211364.5	83629745.46	23633031.35	0.121056762	0.303448276	0.231014719	0.124617263	0.02015837	169	3				
126	E126	533	112773465.2	16045258.44	1516	520991907.5	17216278.56	409188652.7	0.465290607	0.25	0.374693999	0.131771586	0.071031953	302	4				
127	E127	1460	1646702.46	118856.48	4026	651037664.5	16570221	869742616.6	0.388865517	0.467968271	0.430278466	0.027070227	0.01090969	237	4				
128	E128	3222	8628322.93	435467.02	1255	242390151.6	8290697	241526718.8	0.307572936	0.900876494	0.726830903	0.040652222	0.003350458	80	7				
129	E129	6320	70126149.46	945167.36	5817	324006332.6	4663339.68	288821285.5	0.87891242	0.04046336	0.0856989	0.050151639	0.008470177	247	4				
130	E130	3261	72445153.24	4061822.88	1382	110036672.9	10107749.25	42737448.01	0.207628422	0.923296666	0.609175481	0.040794302	0.008417872	35	7				
131	E131	7077	111348920.1	13756660.1	3336	209004031.8	32676655.58	116656927.1	0.172389431	0.056354916	0.128245843	0.075770671	0.016517814	250	2				
132	E132	9159	64001771.73	6110116.07	1798	20338281.9	17607972.88	150854346.9	0.121444272	0.276326265	0.218613789	0.036471449	0.00567322	218	2				
133	E133	1645	30108647.88	1550384.16	170	113655014.4	340650.48	76405632.82	0.227355623	0.894117647	0.652356094	0.041873276	0.002754821	34	7				
134	E134	1695	3628745.5	366483.89	3070	122815470	7368920.27	125689170.9	0.161061947	0.08456026	0.127306568	0.055623715	0.013641133	297	1				
135	E135	2891	24665279.45	1815677.99	413	135337558.5	7758321.61	116914622.7	0.65863323	0.258079933	0.258079933	0.433628837	0.027230709	0.005768065	183	7			
136	E136	1398	114620006.3	17660033.89	5365	116652185.4	17813209.64	2465215.75	0.80757764	0.976336793	0.86346635	0.027623129	0.00189179	21	0				
137	E137	3281	103200564.5	5032857.25	254	143264372.9	10273700.18	45236531.32	0.185306557	0.496062092	0.374444656	0.046110325	0.012164074	133	2				

Figure 9

Result4 data (part)

coefficient method to test the number of classification is that the closer the curvature of K value is to 1, the higher the accuracy of K value is. As can be seen from Figure 7, for the number of data classification of result3, the curvature of K from 1 to 9 is less than 0. When k = 10, the curvature is close to 1. When k is from 11 to 11, the curvature is less than 1, so the best clustering number should be 10 [12].As a single test method is accidental, we adopt two different test methods, namely elbow method and contour coefficient method, to test the best number of clusters for the same data at the same time, and the best number of clusters is 10.As shown in Figure 8, it is the number of corresponding enterprises that are divided into 10 categories by K-means clustering algorithm.As shown in Figure 9, it is the data result result4 after the algorithm processing, and there is a detailed analysis later.

5.3.2 credit strategy

According to the cluster analysis in 5.3.1, we divide enterprises into 10 categories according to the credit risk from small to large. The specific categories and relative lending strategies are shown in Table 2[13].

Table 2
Divides enterprises into ten categories according to the degree of risk

category	Enterprise code
0	X136 X168 X297 X329 X336 X343 X349 X355 X356 X357 X359 X360 X366 X370 X374 X375 X376 X382 X383 X387 X389 X391 X400 X402 X404 X408 X413 X414 X415 X416 X417 X421 X422 X423 X425
8	X166 X199 X358 X380
6	X192 X210 X215 X226 X230 X231 X232 X233 X236 X237 X240 X241 X243 X244 X245 X247 X249 X254 X255 X263 X266 X267 X268 X269 X270 X274 X277 X279 X291 X292 X293 X296 X303 X305 X307 X313 X314 X316 X318 X320 X321 X322 X324 X326 X327 X328 X331 X332 X333 X334 X335 X338 X340 X341 X342 X351 X354 X361 X362 X365 X367 X368 X369 X371 X372 X373 X377 X378 X381 X384 X386 X388 X390 X393 X394 X396 X398 X399 X403 X406 X407 X409 X410 X412 X418 X419 X424
7	X128 X130 X133 X135 X138 X139 X140 X141 X142 X143 X144 X149 X151 X152 X153 X156 X157 X15 X161 X163 X170 X171 X172 X176 X183 X185 X186 X190 X193 X194 X202 X203 X205 X207 X216 X220 X225 X229 X238 X273
1	X134 X148 X179 X184 X187 X189 X198 X200 X201 X211 X212 X213 X214 X218 X219 X221 X223 X227 X228 X234 X235 X239 X246 X248 X250 X251 X252 X253 X256 X257 X258 X259 X260 X261 X262 X265 X271 X272 X275 X276 X278 X280 X281 X283 X284 X285 X286 X287 X288 X289 X290 X294 X295 X298 X299 X300 X301 X302 X304 X306 X308 X309 X310 X311 X315 X317 X319 X323 X325 X330 X337 X339 X344 X345 X346 X347 X348 X350 X352 X353 X363 X364 X379 X392 X395 X411 X420
3	X124 X125
2	X131 X132 X137 X145 X146 X147 X150 X154 X155 X159 X160 X162 X164 X165 X167 X169 X173 X174 X175 X177 X178 X180 X181 X182 X188 X191 X195 X196 X197 X204 X206 X208 X209 X222 X224 X312
9	X139 X141 X171 X185 X193 X194 X225 X238 X273
5	X217 X242 X264 X282 X385 X397 X401 X405
4	X126 X127 X129

According to the degree of risk, it is divided into ten categories. Banks should give priority to the ten categories and gradually reduce the amount of loans. When the total amount of loans is 100 million yuan, the enterprises in the front of the ladder give priority to the amount of loans. When the total amount of loans of all the enterprises in front exceeds 100 million, the enterprises in the back step will not be able to obtain loans. The loan

Table 3
Loan scheme

Types (from low risk to high risk)	Enterprise loan amount (million)	Proportion of bank loan amount (%)
0	[0,100]	100
8	[0,90]	95
6	[0,80]	90
7	[0,70]	85
1	[0,60]	80
3	[0,50]	75
2	[0,40]	70
9	[0,30]	65
5	[0,20]	60
4	[0,10]	55

amount and proportion of different types of enterprises are different, and the specific loan scheme is shown in Table 3. The interest rate of each enterprise has been calculated, as shown in the appendix material result3 data.

6. Conclusions and suggestions

6.1 Model Evaluation

The model established by qualitative and quantitative analysis of corporate credit risks is more rigorous and comprehensive in scope. The model avoids the subjectivity of data, does not rely on objective functions or test results, and accurately describes the overall impact of influencing indicators. In addition, there are no strict limitations on sample size, data distribution, and number of indicators, making it suitable for small sample data and large-scale systems with multiple units evaluation and indicators, and more convenient and flexible.

6.2 Future improvement

The solutions and methodologies are extensively applied in various field [14-17] with effectiveness and beneficial. This solution will be considered by the following items: (1) DNN dense link NN algorithm is easy to over fit and fall into optimum local points. (2) It is difficult to select the corresponding quantitative indicators for the data of each indicator. (3)

It is impossible to predict the credit rating of enterprises in different periods.

Acknowledgements

The authors are appreciative to the study grants from Yahui Meng at the Provincial platforms key and major research scientific universities projects at Guangdong Province, P.RC under Grant Number 2017GXJK116 and Guangdong Provincial Higher Education Association Laboratory Management Professional Committee Foundation, Peoples R China under Grant No. GDJ2016048. These authors were grateful to the grants research given to ZY Chen from the Projects of Talents Recruitment of GDUPT (NO. 2021rc002) in Guangdong Province, Peoples R China, and the research grants given to Ruei-Yuan Wang from the Projects of Talents Recruitment of GDUPT, Peoples R China under Grant NO. 2019rc098 as well as to the anonymous reviewers for constructive suggestions.

References

- [1] A. E. Tarabishy, "The Genesis of the United Nations International Name Day for Micro-, Small, and Medium-Sized Enterprises — June 27," *J. Int. Council for Small Bus.*, vol. 1, pp. 4–6 (2020).
- [2] M. Bhattacharya, S. R. Paramati, I. Ozturk, and S. Bhattacharya, "The effect of renewable energy consumption on economic growth: Evidence from top 38 countries," *Appl. Energy*, vol. 162, pp. 733–741 (2016).
- [3] S. Özokcu and O. Ozdemir, "Economic growth, energy, and environmental Kuznets curve," *Renewable Sustainable Energy Rev.*, vol. 72, pp. 639–647 (2017).
- [4] E. Oduro-Ofori, P. Anokye, and M. Edetor, "Microfinance and Small Loans Centre (MASLOC) as a model for promoting micro and small enterprises (MSEs) in the Ashaiman Municipality of Ghana," *J. Econ. Sustainable Dev.*, vol. 5, pp. 53–65 (2014).
- [5] N. Yoshino and F. Taghizadeh-Hesary, "Optimal credit guarantee ratio for small and medium-sized enterprises' financing: Evidence from Asia," *Econ. Anal. Policy*, vol. 62, pp. 342–356 (2019).
- [6] J. Guo, X. Liu, C. Cui, and F. Gu, "Influence of nonspecific factors on the interest rate of online peer-to-peer microloans in China," *Finance Res. Lett.*, no. 101839 (2020).

- [7] T. Havránek, Z. Irsova, and J. Lešánovská, "Bank efficiency and interest rate pass-through: Evidence from Czech loan products," *Economic Modelling*, vol. 54, pp. 153-169 (2016).
- [8] Y. Lin and T. Chen, "A multibelief analytic hierarchy process and nonlinear programming approach for diversifying product designs: Smart backpack design as an example," *Proceedings of the Institution of Mechanical Engineers, Part B: Journal of Engineering Manufacture*, vol. 234, pp. 1044-1056 (2020).
- [9] X. Chenghua, "AHP and its application," *Journal of Lanzhou Business University*, no. 02, pp. 79-82 (2001).
- [10] X. Yang, Y. Cai, J. Tan, and S. Zhang, "Evaluation of economic resilience level in Western China and its influencing factors," *Journal of Hubei Normal University (Philosophy and Social Sciences Edition)*, vol. 41, no. 03, pp. 25-31 (2021).
- [11] H. Asri, H. Mousannif, and H. A. Moatassime, "Reality mining and predictive analytics for building smart applications," *Journal of Big Data*, vol. 6, pp. 1-25 (2019).
- [12] X. Ma, G. Duan, J. Wang, and H. Xue, "Research on airline customer grouping based on improved contour coefficient method," *Operations Research and Management*, vol. 30, no. 01, pp. 140-146 (2021).
- [13] C. Han, "Optimization of K-means clustering algorithm and its application in e-commerce platform precision marketing," *Dissertation, Shandong University of Science and Technology* (2020).
- [14] Z. Chen, "Interconnected TS fuzzy technique for nonlinear time-delay structural systems," *Nonlinear Dynamics*, vol. 76, pp. 13-22 (2014). <https://doi.org/10.1007/s11071-013-0841-8>.
- [15] Z. Chen, "A criterion of robustness intelligent nonlinear control for multiple time-delay systems based on fuzzy Lyapunov methods," *Nonlinear Dynamics*, vol. 76, pp. 23-31 (2014). <https://doi.org/10.1007/s11071-013-0869-9>.
- [16] K. M. Shafi, M. R. Bhat, and T. A. Lone, "Sentiment analysis of print media coverage using deep neural networking," *Journal of Statistics and Management Systems*, vol. 21, no. 4, pp. 519-527 (2018). <https://doi.org/10.1080/09720510.2018.1471263>.
- [17] S. Gupta, L. Vijayvargy, and K. Gupta, "Assessment of stress level in urban areas during COVID-19 outbreak using critic and topsis: A case of Indian cities," *Journal of Statistics and Management Systems*, vol. 24, no. 2, pp. 411-433 (2021). <https://doi.org/10.1080/09720510.2021.1879470>.