

Automatic Business Location Selection through Particle Swarm Optimization and Neural Network

Yinyi Luo

South China University of Technology

Jinghui Zhong (✉ jinghuizhong@scut.edu.cn)

South China University of Technology

Wei-Li Liu

Guangdong Polytechnic Normal University

Wei-Neng Chen

South China University of Technology

Research Article

Keywords: Business location selection, Machine learning, Adam optimizer, Comprehensive learning particle swarm algorithm

Posted Date: July 29th, 2022

DOI: <https://doi.org/10.21203/rs.3.rs-1870959/v1>

License: © ⓘ This work is licensed under a Creative Commons Attribution 4.0 International License.
[Read Full License](#)

Automatic Business Location Selection through Particle Swarm Optimization and Neural Network

Yinyi Luo¹, Jinghui Zhong^{1*}, Wei-Li Liu² and Wei-Neng
Chen¹

^{1*}School of Computer Science and Engineering, South China
University of Technology, Panyu, Guangzhou, 510641,
Guangdong, China.

²School of Computer Science, Guangdong Polytechnic Normal
University, No.293 West Zhongshan Road, Guangzhou, 510665,
Guangdong, China.

*Corresponding author(s). E-mail(s): jinghuizhong@scut.edu.cn;
Contributing authors: Katherine6Luo@gmail.com;
liuweili@gpnu.edu.cn; cschenwn@scut.edu.cn;

Abstract

Location selection is an important part of running a business. A good business location can greatly increase customer flow, reduce operating costs, and increase business revenue. Therefore, it is necessary to select an appropriate location for business premises. In this research, we propose a method combining machine learning and particle swarm optimization for the business location selection. In the proposed method, we first apply the particle swarm algorithm to optimize the proportional coefficient of the influence of different surrounding commercial places on the site selection of the research site. Then we use a machine learning method to perform the binary classification on the sites to judge whether a geographic location is a good site selection point. The proposed method is compared with several common location selection methods such as random selection method and k-nearest neighbor method. The experimental results have validated the effectiveness of our proposed method for business location selection.

Keywords: Business location selection, Machine learning, Adam optimizer, Comprehensive learning particle swarm algorithm

1 Introduction

Site selection is an important step before the actual operation of an enterprise. The decision-making problem of site selection usually affects the investment income, operating costs, tax policy preferences, sales channels, enterprise competitiveness, enterprise resource utilization and sustainable development of enterprises. Lack of consideration in site selection will lead to problems such as high operating costs, lack of labor, poor sales channels, high logistics costs, and loss of competitive advantages in the industry.

With the transformation and upgrading of the economy, the location requirements of enterprise factories are changing, that is, enterprises are more and more inclined to choose mature industrial parks with perfect industrial environment and supporting facilities, and can provide enterprises with services such as talents, capital, technology, preferential policies and so on.

For the location problem, the traditional solution is based on statistical methods. Methods include web-based GIS [1], method based on Multi Criteria Decision Making process [2], method of building an interactive visual analysis system [3], method of building a simple equilibrium framework [4]. Another kind of site selection method is the heuristic method that is very popular these days. Methods include regression method [5], method of Support Vector Machine [6], method of k-nearest neighbor algorithm [7]. The traditional site selection method has disadvantages such as high trial and error cost, time-consuming and labor-intensive, and current situation lag. These methods can only provide a set of systematic steps to solve the problem without considering the relationship between the global decision factors. For conventional heuristic methods, although methods like k-nearest neighbor in some case are quite suitable, most of the heuristic algorithm is based on the linear decision-making idea. In today's complex and changeable environment, the linear decision-making idea gradually exposes its inherent limitations. The nonlinear decision-making method is the focus and trend of further research.

To overcome the aforementioned limitations, we propose a machine learning method that assigns weights to feature indicators. In our method, we use comprehensive learning particle swarm optimization to calculate the weights of each feature index, and based on these weights, calculate the likelihood that a location information is a good site. Based on the likelihood, we use the location for binary classification as the input of the machine learning binary classification model, and finally use the Adam optimizer to predict whether other coordinate points are good location points based on these input data. Our proposed method has several advantages as follows. First of all, instead of simply considering the direct input of influencing factors as features, we assign weights to each influencing factor to maximize the consideration factors and

improve the accuracy of the prediction model. Second, our model employs a comprehensive learning particle swarm algorithm to search a wider space for the optimal values of weights and categorical variables for the distances of various other commercial establishments to maximize its predictive power.

The organization of this paper is as follows. Firstly, it is the literature review on previous studies on site location and the issues around these methods. Then we proposed our method in solving the location selection problem. Finally, we provide experiment studies on our proposed method and compare our method to random selection, greedy algorithm, and k-Nearest Neighbor (NN) to have a general analysis of our proposed method.

2 Preliminary

2.1 Problem definition

Location selection is a very classic problem in operations, which refers to determining the location of the business and rationally planning the business distribution under several given constraints, with the overall goal of optimizing the total service or maximizing social benefits. This paper focuses on the optimization model of the discrete location selection problem and the method to solve the problem.

Business selection can be determined according to the types and distances of the their surrounding buildings. Specifically, given a set of candidate locations for a district and all of its surrounding buildings, choosing a good location can help to maximize the benefits. Since the location selection is a classic NP hard problem, it is rather difficult to obtain the optimal solution in polynomial time.

For the business location selection, surrounding building facilities have different impact on the one we are studying, which together constitute a proportionality factor of 1. Therefore, we use a score value to evaluate how suitable a location to be a business location. The calculating equation of score for a location j is described as follow:

$$Score_j = \frac{1}{\sum_{i=1}^n \alpha_i x_i^2} \quad (1)$$

where α_i is the weight of the i th type surrounding business facility, x_i is the distance between the j th candidate location and its nearest i th type surrounding business facilities. The higher the score value is, the more suitable the location j is to become a business location we are studying.

An algorithm is designed to optimize the weights of different types of the surrounding building facilities, where the weight optimization problem can be described as:

$$\min \sum_{j=1}^m \sum_{i=1}^n \alpha_i x_{i,j}^2 \quad (2)$$

where n is the number of surrounding building facilities that are taken into consideration, m is the number of our candidate location, α_i is the weight of

the i th type surrounding business facility, $x_{i,j}$ is the distance between the j th candidate location and its nearest i th type surrounding business facilities.

2.2 Related work

For the location prediction problem, the first thing to be considered is the factors that will affect site selection. In the literature, the influencing factors of the site selection problem have been widely studied [8]–[9]. Various factors may have a more or less impact on the location selection of commercial premises. People usually consider major factors, such as the equipment of surrounding buildings and facilities [8] and the cost of construction, for location selection.

Based on these factors that have an important impact on the site selection, many site selection methods have been proposed. These methods generally can be divide into two major streams, namely, the traditional method and heuristic method. The traditional solution is based on statistical methods. Chumaidiyah et al. [1] proposed web-based GIS is a user-friendly application to determine the planned industrial location. Md. Al Amin et al.[14] proposed a method based on Multi Criteria Decision Making process for location selection. Kirilin Li et al.[3] proposed a method of building an interactive visual analysis system with a user-friendly interface for an interactive visual query about complex business and environment information for site location decision-making. Talluri et al.[4] proposed a simple equilibrium framework, solvable by integer programming and estimable from public data for location selection.

Another kind of site selection method is the heuristic method that is very popular these days. Bilen et al.[5] proposed several regression model such as SMORegression (SMOReg), MultilayerPerceptron (MLP), and multivariate Linear Regression (Linear) to solve the business location selection problem. Widaningrum et al.[6] proposed a combination of GIS and Support Vector Machine to predicts the class label to solve the location selection problem. Shihab et al. [10] combine Support Vector Machine, Decision Tree and Logistic Regression for comparative analysis and searched for an algorithm that gives the best result for restaurant business location. Yang Yang et al.[11] designed an automated web GIS application for evaluating potential hotel location. The application uses a set of machine learning algorithms to predict various business success indicators associated with location sites. Dehshib et al.[7] proposed a k-nearest neighbor algorithm for site selection which is quite suitable.

The traditional site selection method has disadvantages such as high trial and error cost, time-consuming and labor-intensive, and current situation lag. These methods can only provide a set of systematic steps to solve the problem without considering the relationship between the global decision factors. For conventional heuristic methods, most of them The heuristic algorithms are based on the linear decision-making idea. In today's complex and changeable environment, the linear decision-making idea gradually exposes its inherent limitations. The nonlinear decision-making method is the focus and trend of further research.

Based on these shortcomings, we propose a new method for location selection. Genetic algorithm can be used to optimize the classification. Salama et al. [12] applied ACO to optimize the number of clusters, the positioning of the clusters, and the choice of classification algorithm to use as the local classifier for each cluster. Albinati et al. [13] uses an ACO algorithm as the supervised learning method in the self-training procedure to generate interpretable rule-based models—used as an ensemble to ensure accurate predictions to accomplish semi-supervised self-training algorithm. Smamanta et al. [14] used PSO algorithm to select input features and the classifiers are trained with a subset of experimental data for known machine conditions and are tested using the remaining data, finally proved the effectiveness of the selected features and classifiers in detection of machine condition. In this paper, we use a comprehensive learning particle swarm algorithm to optimize the weights of various parameters. The score function determined by this method is used as the standard of the binary classification method, and the site selection site is classified into 0 and 1 by the adam optimizer for site selection. In our proposed method, we use the genetic algorithm, a heuristic algorithm with random search, which has strong global search ability. The genetic algorithm is combined with the binary classification algorithm to avoid falling into the local optimal situation, which makes it more valuable for application.

3 Proposed Method

This research is devoted to studying the influence factors of different building types on the location selection, and then the candidate locations are classified as 0 and 1 by the method of machine learning to check whether a candidate location is suitable.

3.1 General Framework of Proposed Method

For better understanding, Fig.1 illustrates the general framework. Coefficient Optimizing Module, Neural Network Training Module and Testing Module are function as optimizing the score calculated coefficient, training the classification neural network and testing the training model separately. To differentiate influence of various types of buildings, a particle swarm algorithm is first used to optimize the weights taken the locations of our research business and the surrounding business as input as it is shown in Coefficient Optimizer Module. Based on the weights that we obtained from the Coefficient Optimizer Module which is X_{best} , the score value of a geographic location can be calculated. Note that score values of the candidate locations of commercial building in a region will be within a relatively high-value range. In other words, a candidate location with a relatively low score value is not likely to be a good one.

According to Eq. (1), we can know that a candidate location with score value is much smaller than that of an existing commercial building is very unlikely to be the same commercial building. Therefore, we believe that a candidate location with a very small score value is not a suitable location

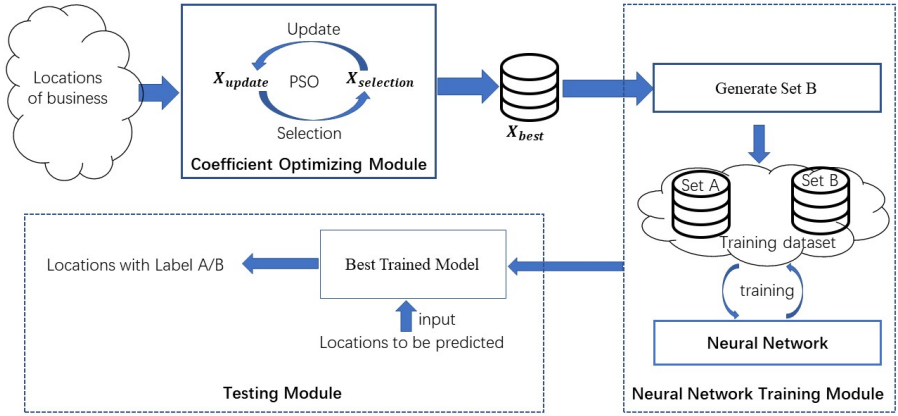


Fig. 1 The framework of our proposed method.

for the commercial building we are studying. To apply the neural network for binary classification, we need to input the existing two types of data sets into the neural network so that the neural network can be trained. The two types of data are in set A and B separately. Locations in set A are locations that are existing suitable locations that we are studying. Locations in set B are those whose score value are at a low level that we believe they are not suitable locations for our research locations. For set A, we use currently existing business location we are studying. For generating set B, we randomly generate some locations, calculate their score value using the coefficient we obtained from Coefficient Optimizing Module, and choose locations where their score value are rather low to set them as type B locations. Taking data in set A and set B as input to the neural network, a well-trained classified module would be produced as it is shown in Neural Network Training Module. Finally, in the Testing Module, enter the locations to be predicted as input into the best trained model, locations with label A or B can be produced. Compare their predicted label and actual label of the input, the accuracy can be calculated.

3.2 Coefficient Optimization Module

As proposed in research [8], the type of building will have an impact on site selection. At the same time, different types of buildings will have different impact factors. Therefore, when studying the impact of different types of buildings on site selection, it should be given different influence values for different types of buildings.

So as to determine the optimal scale of influence of various types of buildings, we used an improved version of the particle swarm algorithm which is the Comprehensive learning particle swarm optimizer[15]. This version of the particle swarm algorithm uses all the best particle information in history to

update the latest individual information, so as to preserve the diversity of the group and prevent problems such as premature convergence, which makes it more effective in solving multi-modal problems than other particle swarm algorithms [16].

Applying CLPSO to get the scale coefficients of the influence of various types of business places, we combine the influence scale coefficients of various business places into a one-dimensional vector as the position vector $x_i = [x_i^1, x_i^2, \dots, x_i^D]$ to be optimized where D is the dimension of the problem and i is the serial number of a particle. The fundamental operations performed by CLPSO are initialization, update and selection. These operations are described as follows.

3.2.1 Initialization

In order to initialize the problem, we need to provide the velocity $v_i = [v_i^1, v_i^2, \dots, v_i^D]$ and position $x_i = [x_i^1, x_i^2, \dots, x_i^D]$ for every individual (particle) in the population where D is the dimension of the problem and i is the serial number of a particle. Normally, the initial velocity vector and position vector of the individual are set randomly. Meanwhile, we need to set the individual's historical best **pbest** as the current position, and the best individual in the group as the current **gBest** which is the global optimal individual.

3.2.2 Update

In this operation, we aim to modify the velocity vector and position vector to search for the optimal solution for the problem. For each particle, randomly select two particles in the entire population exclude the particle which has been updated, compare the fitness of the two particles and take the better one as the candidate. Then, according to the cross probability P_C , the particle to be selected is crossed with the historical probability of the particle according to the dimension to generate a comprehensive learning factor. It means that we randomly generate a number between 0 and 1. If this number is greater than P_C , then we let this dimensional vector of this particle learn from its own **pbest**, otherwise, learn from the **pbest** of the particles in the candidate population. The update algorithm is written as

$$V_i^d = w * V_i^d + c * rand_i^d * (pbest(f_i(d))^d - X_i^d) \quad (3)$$

Where w is the inertia weight, V_i^d is the velocity of the i th particle on the d th dimension, c is the acceleration constant, $rand_i^d$ is a random number for the d th dimension, $pbest(f_i(d))^d$ is the best previous position yielding the best fitness value for the i th particle, X_i^d is the velocity value of the i th particle in the d th dimension. The candidate population of a particle's learning is always unchanged, unless after learning from this population, the number of times the particle's fitness function does not improve beyond the threshold, then the learning population will be updated.

3.2.3 Selection

Machine learning algorithms can be trained based on a large amount of historical data to find certain patterns and make predictions. Classification problem is a core problem of supervised learning, which learns a classification decision function or classification model (classifier) from data, makes output predictions for new input, and the output variable takes a finite number of discrete values[11]. In our proposed method, we use a binary classification method[10] to predict geographic locations to determine whether a geographic location is a good location for business establishments. The optimizer we use in solving this classification is Adam[12].

The converged optimal position vector obtained after many iterations is the scale factor we are searching for. After a certain number of iterations, the population will converge. At this point, the **gbest** individual is the solution we want. In the location selection problem, suppose that we want to predict the location of A, we take the average of the squares of the distances of each McDonald's on the map from the other five buildings closest to it, as the fitness function of CLPSO to optimize the optimal influence ratio. In other words, suppose that $d_{i1}, d_{i2}, d_{i3}, d_{i4}, d_{i5}$ are the closest distance for the i th building A to reach the j th type of other building type, the equation of the fitness function would be

$$\phi = \frac{1}{\sum_{i=1, j=1}^{n, m} x_j * d_{ij}^2} \quad (4)$$

where x_j is the j th building's influence ratio.

3.3 Neural Network Training Module

Neural Networks can be trained based on a large amount of historical data to find certain patterns and make predictions. For the matter of fact, neural network can be used to train classification model. A neural network learns a classification decision function or classification model (classifier) from data, makes output predictions for new input, and the output variable takes a finite number of discrete values [17]. In our proposed method, we use neural network in sequential model [18], a very useful classification model to predict geographic locations to determine whether a geographic location is a good location for business establishments.

Sequential models are linear stacks of layers where one layers leads to the next. The deep neural network is built by stacking many layers. The data is trained layer by layer in sequence where the previous layer is the input to the next layer. In this model, we utilize plain stack of layers where each layer has one input and one output tensor. Each one of these layers and nodes in the model is like a tensor. In sequential model, it does not allow create models that share layers or have multiple inputs or outputs. For time series in sequential model, it just goes from one input to the other while producing an output at each time step during.

The operation of this model includes defining the model, defining the optimization objective, inputting data, training the model, and finally evaluating

the performance of the model using test data. The first layer of the network is the input layer, we need to input training data to the network. To achieve the goal, we need to specify the size of the training data provided to the network, and specify the shape of the input data. The next step is to add intermediate layers and output layers to the model, which is similar to the approach of adding the input layer. Then, it would be the most important part, which is choosing the optimizer and specifying the loss function to specify how backpropagation is calculated. The optimizer we use in solving this classification is Adam [19]. Adam optimizer involves a combination of two gradient descent methodologies. In our loss function we use categorical crossentropy [20], which is used to evaluate the difference between the probability distribution obtained by the current training and the true distribution. It depicts the distance between the actual output (probability) and the expected output (probability), that is, the smaller the value of cross entropy, the closer the two probability distributions are. The characteristic of this kind of loss function is that it is updated according to the error, so when the error is large, the weight update is fast, and when the error is small, the weight update is slow.

4 EXPERIMENT STUDIES

In this section, we apply the proposed method to real cases to test the effect of our proposed method. The test dataset is the geographical distribution of various commercial buildings in Tianhe District and Yuexiu District, Guangzhou City, Guangdong Province, China. We selected McDonald's locations as the object of our study and applied our proposed method to analyze the prediction results. For each example, we compare our proposed algorithm with random selection, greedy algorithm, and k-Nearest Neighbor (NN).

4.1 Experiment Settings

4.1.1 Test Scenarios

Case1: Tianhe District

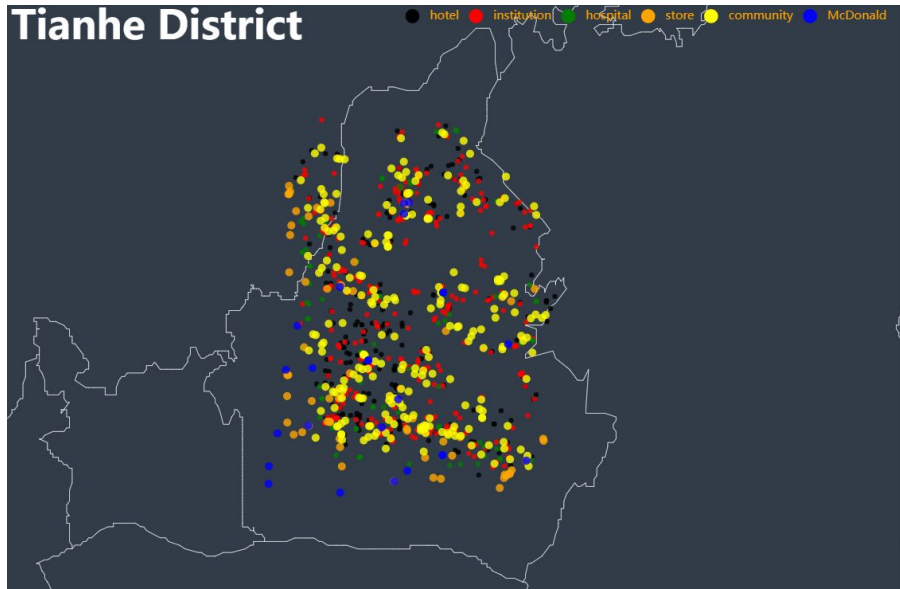
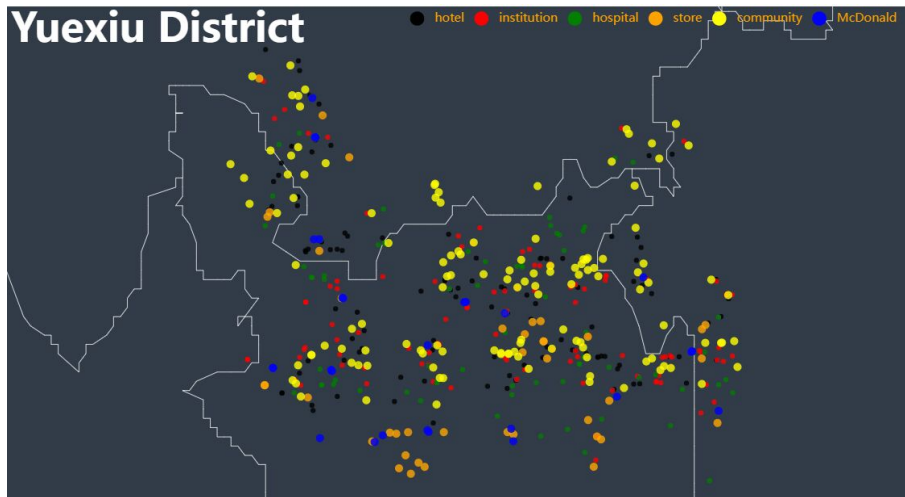
In this case study, our data includes the location latitude and longitude of six places in Tianhe District: 'education education institution', 'hotel', 'store', 'community', 'hospital', and 'McDonald'. The distribution of the six places in Tianhe District are shown in Figure 2.

Case2: Yuexiu District

In this case study, our testing data includes latitude and longitude of the six places in Yuexiu District same as Tianhe District, which is 'education institution', 'hotel', 'store', 'community', 'hospital', and 'McDonald'. The distribution of the six places in Yuexiu District are shown in Figure 3.

4.1.2 Comparison Algorithms

We will compare our proposed method with three different algorithms which are random selection, greedy algorithm and K nearest neighbor method. All

**Fig. 2** Places Distribution in Tianhe District**Fig. 3** Places Distribution in Yuexiu District

the methods selected for comparison are based on the previous studies on location selection. The first two comparison are aiming to compare our proposed method with the statistical methods, and for K nearest neighbor method which performs quite well on location selection problem as we mentioned above, we are aiming to compare our algorithm with heuristic method.

Random Selection

For random selection, the probability of judging whether a location is suitable is 50% which means the accuracy of random selection method is 50%. Compare with random selection is the same as to evaluate the accuracy is larger or smaller than 50%.

Greedy Algorithm

Applied to the above study case, the greedy algorithm is designed as the distance between the selected point and the nearest five places is as small as possible. Suppose that $x_{1i}, x_{2i}, x_{3i}, x_{4i}, x_{5i}$ are the smallest distant from the i th input point to education institution, hotel, store, community and hospital separately. Then the greedy algorithm is aim to choose the points with smallest $\sum_{j=1}^5 x_{ji}$, where the range of i is all the input point. When we apply the greedy algorithm, we choose according to the proportion of the coordinates of McDonald's in the input sample, and the selected proportion that matches the real one is the correct rate of the greedy algorithm. For example, if the percentage of real McDonald location is α . Then for the testing sample, we choose the first $\alpha\%$ of input location with the smallest $\sum_{j=1}^5 x_{ji}$. We use the accuracy calculated in this way to compare with the accuracy calculated by our proposed method.

K Nearest Neighbor Method

KNN is a supervised machine learning algorithm that can be used to solve the classification problem and performs quite well. The KNN algorithm works by assuming that similar things exist in close proximity. The coordinate points on the map have six types of buildings, so in this case, we define the selection method of K Nearest Neighbor as selecting 50 coordinate points closest to the coordinate point. If the coordinate point of McDonald's reaches more than one sixth, the point is considered to be McDonald's, otherwise, the point is considered not to be McDonald's. We take the probability that the true McDonald's coordinates are selected as the correct rate to evaluate the KNN algorithm and compare it with our proposed method.

4.1.3 Performance Metrics and Parameter Settings

To calculate the influence ratio of various types of buildings, we take the average of the squares of the distances of each McDonald's on the map from the other five buildings closest to it, as the fitness function of CLPSO to optimize the optimal influence ratio. In other words, suppose that $d_{i1}, d_{i2}, d_{i3}, d_{i4}, d_{i5}$ are the closest distance for the i th McDonald to reach education institution', 'hotel', 'store', 'community', 'hospital' separately, the fitness function would be

$$\phi = \frac{1}{\sum_{i=1}^n (x_1 * d_{i1}^2 + x_2 * d_{i2}^2 + x_3 * d_{i3}^2 + x_4 * d_{i4}^2 + x_5 * d_{i5}^2)} \quad (5)$$

where x_1, x_2, x_3, x_4, x_5 are the ratio for 'education institution', 'hotel', 'store', 'community', 'hospital' separately. In this way, we obtain the value of x_1, x_2, x_3, x_4, x_5 separately.

Then we randomly select 50 coordinate points on the map of Yuexiu District, calculate the fitness value of these 50 coordinate points, and select 20 coordinate points with the smallest fitness value and far lower than the fitness value of McDonald's places as the point of type 0 (not McDonald's). Then we randomly select 30 McDonald's coordinate points as type 1 (McDonald's coordinate points), and input these two types of coordinate points into the neural network to test the classification result.

Accuracy Measurement For our proposed method, the correct rate is calculated as the proportion of McDonald's coordinate points in the test set that are classified as suitable by the model to all McDonald's in the test set. For the random selection method, since there are only two types of yes or no for binary classification, the correct rate of random selection is close to 50% when the data size is large. Therefore, we use 50% as the correct rate of the random selection method. For the greedy algorithm, we aim to select the coordinate point that is the smallest from the weighted average of various facilities. Therefore, we select the coordinate point with the smallest weighted average in the test set according to the proportion containing McDonald's coordinates in the test set, and define the correct rate as the number of selected points that are actually McDonald's coordinate points in the total number of McDonald's in the test set. For the K nearest neighbor method, we select the 50 points closest to the input point, and if these 50 points are more than one-sixth of the McDonald's coordinate points, the input point is considered as a suitable location point. Therefore, we define the correct rate as the proportion of all McDonald's coordinate points in the test set that are judged as suitable location points.

4.2 Results and Analysis

The accuracy rate of the two study cases using our proposed method and comparison methods are shown in table 1 and we blacked out the best results of these methods.

Table 1 Accuracy of test cases

Method	Accuracy on case1	Accuracy on case2
Proposed method	74%	76%
Random selection	50%	50%
Greedy algorithm	62%	65%
KNN	69%	52%

As can be seen from this table, the accuracy of our proposed method is around 75%. Compared with the random site selection method, it is about 25% higher. The accuracy of the greedy algorithm is about 65%, and the accuracy of our proposed method is about 10% higher than that of the greedy algorithm. The accuracy rate of the K nearest neighbor method is not stable, but its

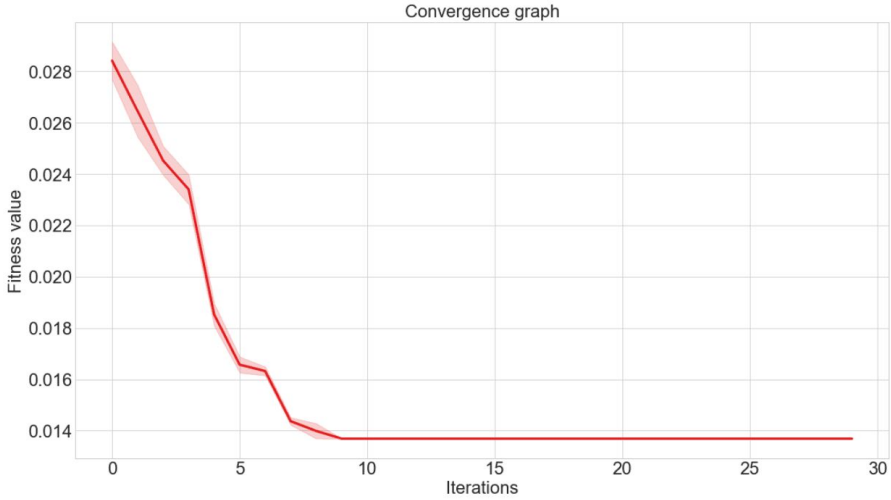


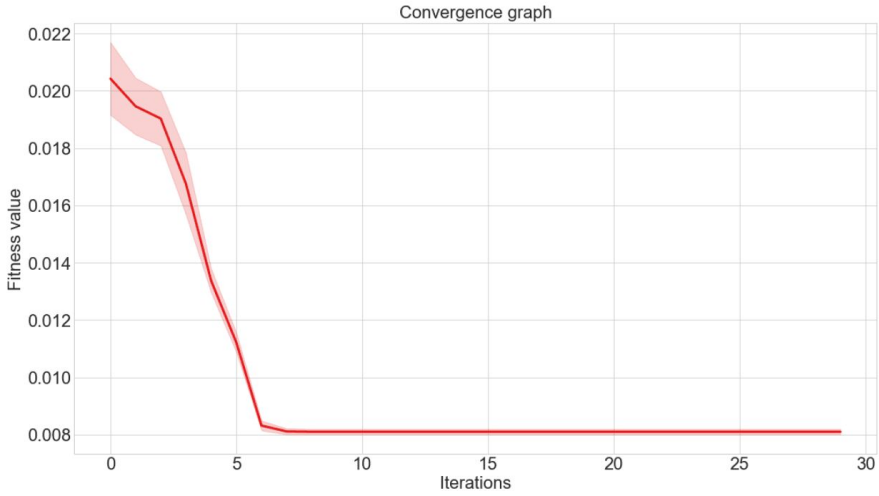
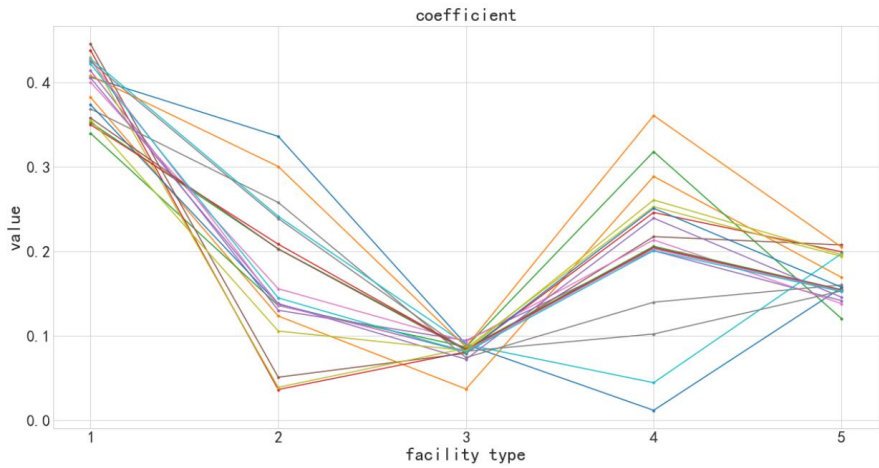
Fig. 4 Convergence curve of case1

highest accuracy rate is about 70%, and our proposed method is still about 5% higher than the highest accuracy rate of the K nearest neighbor method.

We use the CLPSO algorithm to optimize the weight coefficients of various facilities. After running it 20 times, we plot the average value of each generation on the graph, and the shaded area represents the variance. The convergence curve of case 1 and case 2 are shown in Figure 4 and Figure 5 separately. It can be seen that the curves are convergent at approximately the 9th and 7th iteration and reach the global optimum in the two cases. What the two cases have in common is that as the iteration increases, the variance of the fitness value gradually decreases, and when the curve converges, the variance value reaches a minimum. The difference between the two cases is that the reduction rate of the variance and convergence rate of case1 are slower than those of case2. At the beginning, the variance of case 1 is smaller than that of case 2, and it does not converge until the ninth iteration, while the variance of case 2 is larger than that of case 1 at the beginning, and the convergence is basically completed by the seventh iteration.

We repeated the experiment 20 times when using CLPSO to calculate the impact ratios of various facilities. The result of the proportion for ‘education institution’, ‘hotel’, ‘store’, ‘community’, ‘hospital’ in case 1 and case 2 are shown as in Table 2 and 3 and the value are visualized in Figure 6 and Figure 7. One to five in facility type in the two figures correspond to ‘education institution’, ‘hotel’, ‘store’, ‘community’, ‘hospital’ respectively. From the data in the tables and images we can draw some conclusions.

case1 The weight ratio of hotel has always been at a high level, indicating that hotel is a key factor in the selection of McDonald’s in Tianhe District. The proportion and weight of store and hospital fluctuate greatly, indicating that store has no obvious rules for the location of McDonald’s in Tianhe District.

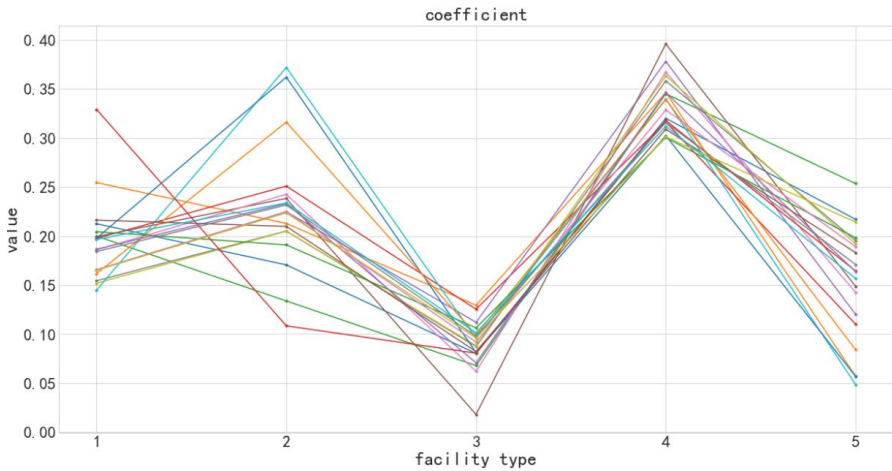
**Fig. 5** Convergence curve of case2**Fig. 6** Coefficient of case1

The community is at a stable and small level, indicating that its impact on the location of McDonald's in Tianhe District is relatively small. The hospital is at a stable and moderate level, indicating that it has a certain influence on the location of McDonald's in Tianhe District.

case2 The influencing factors of community have been kept in a relatively large range, indicating that community has a dominant influence on the location of McDonald's in Yuexiu District. The impact of store has been at a relatively low level, indicating that it is a dispensable impact factor. Education institutions, hotels and hospitals fluctuate greatly, and there is no obvious rule affecting the location of McDonald's in Yuexiu District.

Table 2 Weight of various facilities of case1

Number	education institution	hotel	store	community	hospital
1	0.405881	0.33583941	0.089794	0.01180185	0.15668314
2	0.38204811	0.12331927	0.03701158	0.28848466	0.16913637
3	0.33948977	0.13568437	0.08660472	0.31774214	0.12047901
4	0.4375855	0.03635121	0.08096373	0.24582032	0.19927924
5	0.41338693	0.1301461	0.09433563	0.20058575	0.1415456
6	0.44499152	0.05097735	0.0793147	0.21726426	0.20745217
7	0.39981756	0.15563715	0.09313682	0.21337143	0.13803704
8	0.36836959	0.25763354	0.07419496	0.13968757	0.16011433
9	0.42914877	0.03902448	0.08520164	0.25266225	0.19396286
10	0.42759907	0.24125383	0.08972219	0.04450176	0.19692315
11	0.37327215	0.13747865	0.08096994	0.25075269	0.15752657
12	0.40752374	0.30018561	0.08686562	0.36044523	0.20497979
13	0.35233762	0.20280608	0.08442748	0.20584379	0.15458504
14	0.34964449	0.20830351	0.08349911	0.20454484	0.15400805
15	0.40496041	0.13811019	0.07219523	0.2391173	0.14561686
16	0.35752798	0.20262388	0.08276952	0.20352401	0.1535546
17	0.42586464	0.13601923	0.08219614	0.20272175	0.15319824
18	0.42455743	0.23868759	0.08174553	0.10209127	0.15291818
19	0.35544293	0.10558175	0.08247389	0.26056435	0.19593708
20	0.42158961	0.14474575	0.08072247	0.20065983	0.15228234

**Fig. 7** Coefficient of case2

It can be seen that in the two cases, the proportion of hotels is close to half, indicating that hotels are very critical to the location of McDonald's. The reason for this could be that hotel can lead to tourism-related development and promotion and then gradually promote tourist demand for McDonald's food service. Apart from that, education institution also has a relatively high

Table 3 Weight of various facilities of case2

Number	education institution	hotel	store	community	hospital
1	0.21240688	0.17052931	0.08026905	0.31986985	0.21692491
2	0.25444905	0.21273822	0.12950544	0.34612817	0.05717911
3	0.19985735	0.13379769	0.06792124	0.34494604	0.25347768
4	0.19756385	0.25082984	0.12538317	0.31613113	0.11009202
5	0.1651611	0.22475767	0.11204613	0.37787388	0.12016122
6	0.19931163	0.23822307	0.01815326	0.39561424	0.1486978
7	0.18583637	0.242405	0.06211569	0.36698066	0.14266228
8	0.15437737	0.20503826	0.0872393	0.35810938	0.19523568
9	0.15187402	0.2048633	0.08826022	0.3631081	0.19189437
10	0.14481191	0.37165762	0.09540202	0.33948865	0.0486398
11	0.19769463	0.36186359	0.08125897	0.30216676	0.05701606
12	0.16165729	0.31608324	0.09915241	0.33899899	0.08410807
13	0.20433993	0.19091405	0.10635335	0.30028623	0.19810644
14	0.32906754	0.10837874	0.08053303	0.3176957	0.164325
15	0.18652574	0.23301598	0.07070617	0.34596896	0.16378315
16	0.21619696	0.20973022	0.08226186	0.30883477	0.18297619
17	0.16573427	0.22477947	0.09187904	0.32844826	0.18915897
18	0.1841317	0.23151802	0.09766774	0.31575094	0.17093161
19	0.16512986	0.22359532	0.096075	0.30098377	0.21421605
20	0.1964141	0.23364384	0.10126936	0.3119992	0.1566735

proportion in both cases. The reason for this could be that there is a lot of population movement in educational institutions, and the demand for McDonald's food in the nearby area will also be relatively large.

5 Conclusion

In this paper, we propose a combination of particle swarm algorithm and binary machine learning approach for business location selection. The main idea of the method we propose is to use the particle swarm algorithm to optimize the influence ratio of various types of commercial places, set the fitness function as the total influence, and then determine whether a location is a suitable choice according to the fitness value. This fitness value is served as the input category selection criteria for binary classification. Then, according to the input two types of coordinate information, we use the Adam optimizer to classify the study location points to predict whether the location point is a suitable location point. We applied this method to two test cases and found that the accuracy was around 75%. Apart from that, we compare our proposed method with random selection, greedy algorithm, and KNN method, and find that our proposed method is superior to these algorithms. It can be concluded that the method we propose in this paper can serve as a promising alternative to the existing methods for business location selection.

Declarations

- Funding

Not applicable

- Conflict of interest/Competing interests
All authors disclosed no relevant relationships
- Ethics approval
Not applicable
- Consent to participate
Not applicable
- Consent for publication
Not applicable
- Availability of data and materials
The authors confirm that the data supporting the findings of this study are available within the article and its supplementary materials.
- Authors' contributions
Yinyi Luo: methodology, software, writing.
Jinghui Zhong: Methodology, conceptualization, supervision.
Wei-Li Liu: writing
Wei-Neng Chen: writing and supervision.
All authors reviewed the manuscript.
- Acknowledgements
We would like to thank the anonymous referees for their comments and useful suggestions.

References

- [1] E. Chumaidiyah, M.D.R. Dewantoro, A.A. Kamil, M. Keshavarz-Ghorabae, Design of a participatory web-based geographic information system for determining industrial zones. *Appl. Comp. Intell. Soft Comput.* **2021** (2021). <https://doi.org/10.1155/2021/6665959>
- [2] M. Al Amin, A. Das, S. Roy, M.I. Shikdar, Warehouse selection problem solution by using proper mcdm process. *Int. J. Sci. Qual. Anal* **5**(2), 43–51 (2019)
- [3] K. Li, Y.N. Li, H. Yin, Y. Hu, P. Ye, C. Wang, Visual analysis of retailing store location selection. *Journal of Visualization* **23**(6), 1071–1086 (2020)
- [4] K. Talluri, M. Tekin, Optimal location for competing retail service facilities. Available at SSRN 3583413 (2020)
- [5] T. Bilen, M. Erel-Özçevik, Y. Yaslan, S.F. Oktug, in *2018 IEEE 20th International Conference on High Performance Computing and Communications; IEEE 16th International Conference on Smart City; IEEE 4th International Conference on Data Science and Systems (HPCC/SmartCity/DSS)* (IEEE, 2018), pp. 1314–1321

- [6] D.L. Widaningrum, in *Proceedings of the 9th International Conference on Machine Learning and Computing* (2017), pp. 112–116
- [7] M.M. Dehshibi, M.S. Fahimi, M. Mashhadi, in *Recent Advances in Information and Communication Technology 2015* (Springer, 2015), pp. 133–142
- [8] S.M. Kimelberg, E. Williams, Evaluating the importance of business location factors: The influence of facility type. *Growth and Change* **44**(1), 92–117 (2013)
- [9] K. Park, Identification of site selection factors in the us franchise restaurant industry: an exploratory study. Ph.D. thesis, Virginia Tech (2002)
- [10] I.F. Shihab, M.M. Oishi, S. Islam, K. Banik, H. Arif, in *2018 IEEE Region 10 Humanitarian Technology Conference (R10-HTC)* (2018), pp. 1–5
- [11] Y. Yang, J. Tang, H. Luo, R. Law, Hotel location evaluation: A combination of machine learning tools and web gis. *International Journal of Hospitality Management* **47**, 14–24 (2015). <https://doi.org/https://doi.org/10.1016/j.ijhm.2015.02.008>
- [12] K.M. Salama, A.M. Abdelbar, *Learning cluster-based classification systems with ant colony optimization algorithms*, vol. 11 (Springer, 2017), pp. 211–242
- [13] J. Albinati, S.E. Oliveira, F.E. Otero, G.L. Pappa, An ant colony-based semi-supervised approach for learning classification rules. *Swarm Intelligence* **9**(4), 315–341 (2015)
- [14] B. Smamanta, C. Nataraj, Application of particle swarm optimization and proximal support vector machines for fault detection [j]. *Swarm Intelligence* **3**(4), 303–325 (2009)
- [15] J.J. Liang, A.K. Qin, P.N. Suganthan, S. Baskar, Comprehensive learning particle swarm optimizer for global optimization of multimodal functions. *IEEE transactions on evolutionary computation* **10**(3), 281–295 (2006)
- [16] Y. Cao, H. Zhang, W. Li, M. Zhou, Y. Zhang, W.A. Chaovalitwongse, Comprehensive learning particle swarm optimization algorithm with local search for multimodal functions. *IEEE Transactions on Evolutionary Computation* **23**(4), 718–731 (2018)
- [17] F. Osisanwo, J. Akinsola, O. Awodele, J. Hinmikaiye, O. Olakanmi, J. Akinjobi, Supervised machine learning algorithms: classification and comparison. *International Journal of Computer Trends and Technology*

- (IJCTT) **48**(3), 128–138 (2017)
- [18] Y. Even-Zohar, D. Roth, A sequential model for multi-class classification. arXiv preprint cs/0106044 (2001)
- [19] P. Kingma Diederik, J.B. Adam, A method for stochastic optimization. arXiv preprint arXiv:1412.6980 (2014)
- [20] C.H. Chen, P.H. Lin, J.G. Hsieh, S.L. Cheng, J.H. Jeng, in *2020 3rd IEEE International Conference on Knowledge Innovation and Invention (ICKII)* (2020), pp. 200–203. <https://doi.org/10.1109/ICKII50300.2020.9318835>