



International Journal of Web Engineering and Technology

ISSN online: 1741-9212 - ISSN print: 1476-1289

<https://www.inderscience.com/ijwet>

Application of improved K-means algorithm in the cultivation of creative music talents under the needs of sustainable development and transformation

Peng Li, Zeng Fan

DOI: [10.1504/IJWET.2024.10063583](https://doi.org/10.1504/IJWET.2024.10063583)

Article History:

Received:	05 January 2023
Last revised:	10 August 2023
Accepted:	11 October 2023
Published online:	29 April 2024

Application of improved K-means algorithm in the cultivation of creative music talents under the needs of sustainable development and transformation

Peng Li

Institute of Advanced Studies in Humanities and Social Sciences,
Beijing Normal University,
Zhuhai, Zhuhai, 519087, China
Email: liyoungbird2021@163.com

Zeng Fan*

Institute of Advanced Studies in Humanities and Social Sciences,
City University of Macau (RCAE),
Macao, 999078, China
Email: fan.zeng@yandex.com
*Corresponding author

Abstract: In order to cultivate innovative personnel who are adapted to the development of university education, this paper proposes a K-means clustering algorithm (K-means) based on noise reduction autoencoder for the cultivation of creative music talents and explores the difficulties in cultivating innovative talents. The results show that the research-designed method outperforms K-means on the performance metrics NMI, AMI and FMI for the same dataset. The results of the practical application analysis show that the training of practical operation is weakened in talent training, and the emphasis on practical courses should be strengthened in the subsequent talent training plan.

Keywords: sustainable development; autoencoder; K-means; talent training.

Reference to this paper should be made as follows: Li, P. and Fan, Z. (2024) 'Application of improved K-means algorithm in the cultivation of creative music talents under the needs of sustainable development and transformation', *Int. J. Web Engineering and Technology*, Vol. 19, No. 1, pp.4–19.

Biographical notes: Peng Li receive his Master of Arts (MFA), is currently the Director of Art Education, School of Art and Communication, Beijing Normal University Zhuhai and music teacher of Zhuhai Campus of Beijing Normal University. He has 20 years of experience in vocal music teaching, art talent training and art appreciation teaching. He has published articles in many international and domestic journals.

Zeng Fan received his Doctor of Art, and a teacher of Art Education and Research Center, Zhuhai Campus, Beijing Normal University. He has more than ten years of experience in art teaching and management.

1 Introduction

At present, China has entered a new era of development, and all walks of life have increased demands for industrial upgrading and economic structure adjustment. The new round of education reform also urgently needs to meet the needs of the society and incorporate the concept of sustainable development (Avik et al., 2022). In recent years, the state has made many important arrangements in the field of higher education, and proposed a strategic decision to build a world-class university with Chinese characteristics (Dong and Gao, 2021). With the continuous improvement of teaching management level in universities, how to improve teaching methods and improve teaching content through technological means has become the key. This is of great significance to improve the teaching quality and management level, as well as to improve the learning outcomes of students (Nithyanandam, 2020). With the development of science and technology, people generate large amounts of data in their daily lives. Therefore, data analysis and classification have become the concern of researchers. The field of education is no exception. Therefore, we use data-driven technology as the main research method, combined with a large number of students' performance data, to conduct a comprehensive analysis of students' learning effects, which is of great practical significance (Zhang and Mao, 2021). This study aims to analyse the problems existing in the cultivation of music courses by combining the noise reduction autoencoder with the K-means algorithm. The research content of this study is mainly divided into the following four parts: The second part is a literature review and analysis related to clustering algorithms and teaching; The third part is the model construction of K-means data mining algorithm based on noise reduction Autoencoder; The fourth part is the performance analysis and application effect analysis of the model; Finally, there is a summary discussion on the model and result analysis data. The purpose of this study is to establish an improved K-means data mining model based on denoising autoencoder, which first reduces the dimensionality of a large amount of teaching or training data, and then performs clustering analysis. As autoencoder belongs to feedforward acyclic and unsupervised learning neural network, it has excellent dimension reduction and denoising ability. The K-means algorithm has the characteristics of simple theory and fast clustering speed, which makes it suitable for many research fields. Therefore, the study combines the two to construct a clustering analysis model.

2 Related work

Many scholars have done detailed research on clustering algorithms. A hybrid ant colony optimisation algorithm for image segmentation was studied by El-Khatib et al. (2021). The experimental results show that the algorithm has higher accuracy than Grub cut (El-Khatib et al., 2021). Song et al. (2021) proposed a robust K-means seismic waveform classification algorithm based on Gaussian-weighted sliding window, and the effectiveness of the algorithm was demonstrated by applying it to actual seismic data in the F3 block. In order to improve the operational efficiency and lifetime of wireless sensor networks, Feng and other scholars proposed a targeted clustering method based on K-means algorithm. The experimental results show that compared with the traditional method, the method proposed in this paper has certain advantages, which can provide a

theoretical reference for subsequent related research (Feng et al., 2020). Lin and Li (2021) proposed a short-term photovoltaic power generation forecast method based on the combination of feature selection K-means++ grey correlation analysis and support vector regression. The experimental results show that the prediction accuracy of this model is significantly improved compared with the other five models, which verifies the effectiveness of the method. To further improve the effect of gene module identification, Zhang and other scholars proposed a new gene module identification method by combining the Newman algorithm and K-means algorithm framework in community detection. Experimental results show that the algorithm performs best in terms of error rate, biological significance, and CNN classification metrics (precision, recall, and F-value) (Zhang et al., 2021). The research team of Ramalingam and Thangarajan (2020) proposed a K-implication dynamic grouping method suitable for the dynamic topological properties of VANET. The proposed technique works well in situations involving an advanced number of beams as well as an ambiguous number of groups. Zhang and Peng (2022) used particle swarm and K-means hybrid clustering to segment agricultural product images in the YCbcr colour space. It can effectively reduce the impact of low light and shadows, and achieve a balance between global search ability and convergence speed. The results show that the segmentation stability and accuracy of this method are better than the traditional particle swarm clustering segmentation method.

Many scholars have made rich researches on talent cultivation or music education. Scholars like Cao believe that innovation is the soul of a nation's progress and the first driving force of development. Building an innovative country is the core of the national development strategy, which requires a large number of talents with innovation consciousness and ability. Colleges and universities take on the important task of cultivating innovative talents, and the training mode is related to the quality of innovative talent training (Cao et al., 2021). Scholars such as Zhang and Sun (2021) believe that the integration of production and education has always been an important direction for the development of vocational education in China. In addition, the college should also actively explore new models, new ways and new ideas under the background of the integration of production and education in higher vocational colleges, and take a series of effective measures to practice the cultivation of highly skilled talents in the direction of film and television media. Peng and Zheng (2020) analysed the existing problems in training engineering talents, and discussed the training mode and development ideas of engineering talents under the new situation. It points out their specific measures in professional orientation and training goals, curriculum system and knowledge structure establishment, and practical teaching. Chen et al. (2021) analysed the current problems of postgraduate students majoring in computer science and technology in China, and proposed a new model for cultivating the innovative ability of postgraduate students under the background of science and education integration. Improve the innovative ability of graduates to meet the needs of modern information industry and high-tech development (Chen et al., 2021). Ma and others believe that in the current music education system, Russian music education is in a leading position, compare it with Chinese music education and apply it to the practice of music education, hoping to provide the innovation and optimisation of music education in China. To guide and improve the level of music education in our country (Ma, 2022). Through an examination of interview data, White found that authentic learning has a unique place in the music classroom in high level music courses for NSW high school students. This is manifested in the use of in-depth inquiry-based and student-centred learning tasks such as video

journals, the use of specialised resources and expertise, and in real-world settings, collaborative learning in and out of the classroom (White, 2020). Sang and Xu (2022) use music aesthetic education as a means of continuously improving the music aesthetic quality of the new generation of young people and establishing music aesthetic concepts, which is an important starting point for improving the quality of the whole people. Music teaching in primary and secondary schools should focus on aesthetic education, with a focus on improving the quality of music teaching, and gradually adopt the correct learning methods of art and music aesthetics. The focus is to apply the knowledge learned to daily life and self-improvement and development, in order to achieve the realisation of human society. A beautiful ideal of harmonious and stable development.

In summary, clustering algorithms at home and abroad are mainly used in image processing and error rate analysis, and the analysis of talent training is rarely combined with algorithms. Therefore, this research uses the noise reduction autoencoder as a dimensionality reduction method for high-dimensional data, and combines the K-means algorithm to perform cluster analysis on the performance data of musical talents, so as to make reasonable suggestions for talent training.

3 Improved K-means algorithm based on noise reduction autoencoder

In order to achieve clustering analysis of data generated in the teaching process of music students, this study uses an automatic encoder to denoise massive data to achieve low-dimensional data representation of high-dimensional data. Then, the K-means algorithm was introduced to mine the dimensionality reduced data for further analysis of talent cultivation data.

3.1 High-dimensional data dimensionality reduction model based on denoising autoencoder

In recent years, AutoEncoder (AE) has attracted the attention of many scholars. AE belongs to the type of feed-forward acyclic and unsupervised learning methods (Hou et al., 2021). It is characterised by self-learning from massive data and harvesting feature information hidden in a large amount of data. The traditional AE consists of two parts: encoder and decoder, corresponding to mathematical expressions such as equation (1) and equation (2).

$$h_1 = S(W_1x + b_1) \quad (1)$$

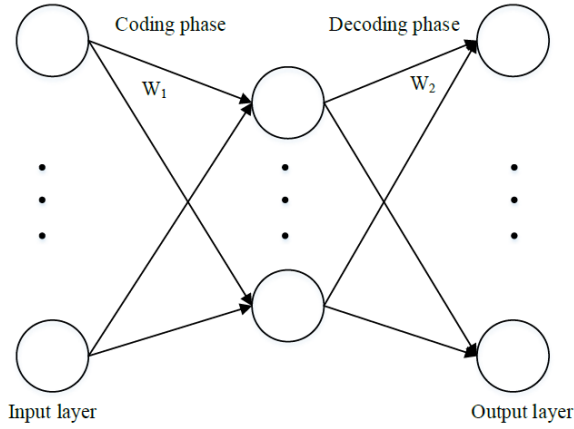
$$y = S(W_2h_1 + b_2) \quad (2)$$

In equations (1) and (2), x and y represent the feature information, W_1 and W_2 represent the weight coefficients of the encoder and the decoder, respectively, b_1 and b_2 represent the partial positive vector of the encoder and the decoder. S represents the activation function during encoding and decoding. h_1 represents the output result of the encoder. Usually, the activation is selected as a sigmoid function, and its equation is as in equation (3).

$$S = \text{sigmoid}(y) \frac{1}{1 + e^{-y}} \quad (3)$$

Figure 1 shows the structural form of the traditional self-encoder. It can be seen that for the self-encoder, the number of nodes in the input layer is the same as that in the output layer. The purpose of network training is to obtain an approximate identity function, and then through the hidden contains layers to capture the necessary information of high-dimensional data, and finally to achieve the effect of using low dimensional data to represent or replace high-dimensional data. In short, the information expressed by the low-dimensional data output by the output layer is y , which is equivalent to the information expressed by the high-dimensional data x .

Figure 1 Structure diagram of traditional self encoder



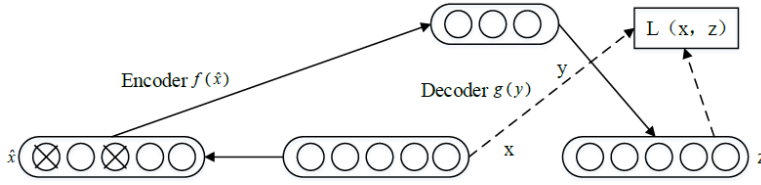
From Figure 1, it can be seen that the purpose of the automatic training of the autoencoder is to make the expression x of y the feature information consistent with the sum. So the network needs to introduce a loss function and at the same time minimise the error between the sum x of y . The mathematical expression of the loss function is shown in equation (4).

$$J_{AE}(W, b) = \sum (L(x, y)) = \sum \|y - x\|^2 \quad (4)$$

In equation (4), (W, b) respectively represent the weight coefficient and the partial positive vector matrix. Although the traditional AE trains the encoder and the decoder separately, in fact, the two neural networks can be connected as an artificial neural network. One is the network training where the input layer points to the hidden layer, and the other is the hidden layer. Training the network pointing to the output layer. The advantage of this approach is that the weight matrix of the encoder and the weight matrix of the decoder are transposed matrices of each other (Ghazal, 2021). At this time, an AE with bound weights is formed, and training the encoder with bound weights together with the decoder can better avoid the negative phenomenon of network overfitting, thus improving the generalisation ability of the network. The traditional AE requires too strict acquisition of identity function, and the network efficiency is too low. Therefore, it is necessary to study more information expression in hidden layers in practical applications,

that is, to ensure that the input and output are approximately the same. In the case of, it is expected to replace the expression of high-dimensional data with data of the lowest dimension, and the higher the sufficiency of the effect of low dimensional expression of high-dimensional data, the better. Unfortunately, the hidden layer of traditional Autoencoder is difficult to reduce high-dimensional data to the extreme when dealing with datasets with high complexity and Big data. It is better to retain the characteristics of the original data (Cui, 2020). To meet the processing requirements of complex data, denoising auto encoder (DAE) emerged as a more robust autoencoder. Figure 2 shows a schematic diagram of the structure of the noise reduction encoder.

Figure 2 Structure diagram of noise reduction encoder



As can be seen from Figure 2, the network first damages the original high dimensional data matrix x to obtain a new matrix, and then $\hat{x}$ pushes the new matrix into the encoder $f(\hat{x})$ as the input of the noise reduction AE. After network training and coding, a low-dimensional data matrix y can be formed. At this time, the role of the decoder $g(y)$ is to reconstruct the low-dimensional data, aiming to map the hidden layer data back z . Its mathematical expression is shown in equation (5).

$$\begin{aligned} y &= f(\hat{x}) = S(W\hat{x} + b_y + b_n) \\ z &= g(y) = S(W'y + b_z) \end{aligned} \quad (5)$$

In equation (5), b_y and b_z represent the partial positive vector of the dimensionality reduction autoencoder. b_n represents random Gaussian noise, and W' is the transpose of W . S is the activation function, which W is taken as the sigmoid function. For the noise reduction autoencoder, the training process is the optimisation process, and the optimised parameters are (W, b_y, b_z) . The corresponding mathematical expression of the reconstruction error is shown in equation (6).

$$J_{DAE} = \sum_{x \in D} L(x, g(f(\hat{x}))) \quad (6)$$

In equation (6), L represents the reconstruction error function. The loss function of the DAE is selected as the cross-entropy loss function, and its mathematical expression is equation (7).

$$L(x, z) = -\frac{1}{n} \sum_{i=1}^{d_x} x_i \ln z_i + (1 - x_i)(1 - z_i) \quad (7)$$

In equation (7), n represents the number of training set samples; x_i represents the i^{th} input data, and z_i represents the i^{th} decoded and reconstructed data. From the above, it can be seen that the denoising autoencoder can pay attention to the role of the hidden layer by

artificially destroying the input data, and allow the neurons in the hidden layer to learn more robust feature expressions.

3.2 *K-means algorithm data mining model optimised based on noise reduction autoencoder*

K-means algorithm has the characteristics of simple theory and fast clustering speed, which is suitable for many research fields. K-means reduces the value of the loss function through dynamic clustering and then iteratively to obtain the optimal solution for clustering. The basic idea of this algorithm is to randomly select data points from the initial data to serve as clustering centres, and then assign other data points to the nearest clustering centre to form many small datasets. Then, the mean of each of these datasets is calculated to obtain a new clustering centre. Finally, iterate repeatedly until the algorithm completion condition is met, and then stop. Therefore, it can be seen that distance calculation is the key element of the K-means algorithm (Anggarwati et al., 2021). Commonly used distance calculation methods include the Euclidean distance method, the Chebyshev distance method, and the Manhattan distance method. The corresponding mathematical expressions are shown in equations (7), (8), and (9).

$$dis(x, y) = \sqrt{(x_1 - y_1)^2 + (x_2 - y_2)^2 + \dots + (x_i - y_i)^2 + \dots + (x_n - y_n)^2} \quad (8)$$

$$dis(x, y) = \max(abs(x_1 - x_2), abs(y_1 - y_2)) \quad (9)$$

$$dis(x, y) = abs(x_1 - x_2) + abs(y_1 - y_2) \quad (10)$$

In equations (8), (9) and (10), $dis(x, y)$ denotes the distance between two data points. x_i and y_i denote the coordinates of the data points, respectively. The Euclidean distance is derived by calculating the distance between two points in Euclidean space, which is expressed as the distance between two points in low dimensional space. Equation (8) represents the Euclidean distance between two data points in n-dimensional space. The Chebyshev distance is usually only suitable for some special cases, and the generality is very low compared to the Euclidean distance method. The Manhattan distance, also known as the checkerboard distance, is also less general. Therefore, the Euclidean distance is the most versatile among the three distance calculation methods, and this study also selects the Euclidean distance as the distance metric. The K-means algorithm includes the expected final number of clusters and the sample data to be clustered when the data is input. Assuming that the sample dataset is $X = \{x^{(1)}, x^{(2)}, \dots, x^{(m)}\}$ the number of clusters, that is, the cluster centre $\{\mu_1, \mu_2, \dots, \mu_k\}$, is selected as k , respectively. Calculate the distance from each data point to the cluster centre according to the Euclidean distance method, as in equation (11).

$$c^{(i)} = \arg \min \|x^{(i)} - \mu_j\|^2 \quad (11)$$

In equation (11), $x^{(i)}$ represents the data target point. $c^{(i)}$ represents the cluster centre closest to the target point, and μ_j represents the j^{th} cluster centre. For each cluster centre, its calculation method is equation (12).

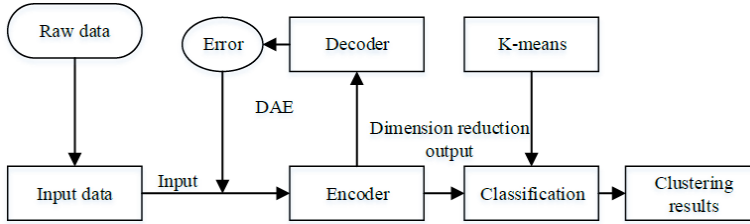
$$\mu_j = \frac{\sum_{i=1}^m 1\{c^{(i)} = j\} x^{(i)}}{\sum_{i=1}^m 1\{c^{(i)} = j\}} \quad (12)$$

The K-means algorithm aims to reduce the sum of the distances between all the sample data points and the cluster centres through repeated calculations and iterations of equation (11) and equation (12). The loss function of the K-means algorithm can be formed by the sum of the squares of the distance differences, and the expression is shown in equation (13).

$$J(c^{(1)}, \dots, c^{(m)}, \mu_1, \dots, \mu_k) = \frac{1}{m} \sum_{i=1}^m \|X^{(i)} - \mu_{c^{(i)}}\|^2 \quad (13)$$

In equation (13), $\mu_{c^{(i)}}$ represents the cluster centre point closest to the sample data point. Based on the loss function, the optimisation goal of the algorithm is to find the sum $\mu_1, \dots, \mu_k$ that minimises the value of the loss function $c^{(1)}, \dots, c^{(m)}$. Considering the fact that the data is usually diverse and the amount of data is huge, the noise reduction AE is used as a data dimension reduction method to combine with the K-means algorithm to achieve the improvement and optimisation of the latter. For the convenience of description, this algorithm is used. It is referred to as the DAE-K algorithm for short. Figure 3 is a schematic flowchart of an improved K-means algorithm based on a noise reduction autoencoder.

Figure 3 Flow chart of improved k-means algorithm for noise reduction self encoder

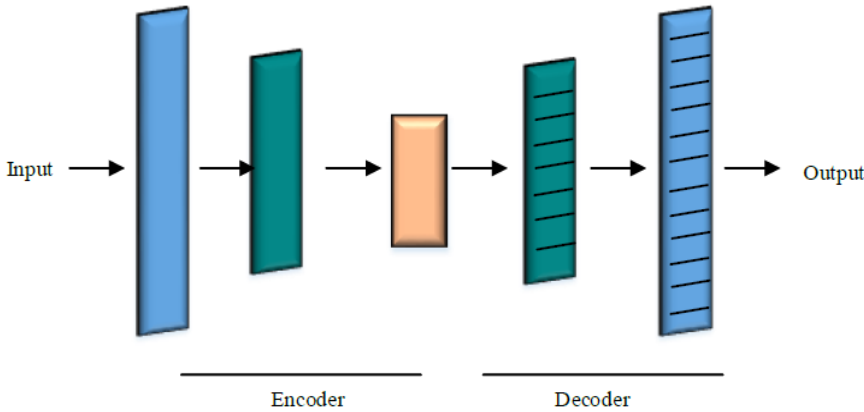


As shown in Figure 3, the DAE-K algorithm consists of two parts, the original data dimension reduction part and the K-means algorithm clustering analysis part. Through the processing of the DAE, the high-dimensional dataset can be reduced to low dimensional data while retaining the original data feature information. Therefore, the dimensionality reduction effect directly affects the clustering effect of the second part of K-means. To improve the DAE, the dimensionality reduction effect of this research also uses a stacked DAE. The schematic diagram of the stacked denoising autoencoder is shown in Figure 4.

As shown in Figure 4, compared with the ordinary denoising autoencoder, the stacked denoising autoencoder is composed of multiple denoising autoencoders and has a larger number of hidden layers. The dimensionality reduction is divided into multiple processes, so that the difference between the original data and the decoded output data is smaller. Therefore, the stacked noise reduction AE can be equivalent to several ordinary noise reduction AEs that reduce dimensionality and stack encoder inputs together. The low-dimensional output data of the previous AE is used as the input data of the next

autoencoder. During decoding, the high-dimensional output data of the previous decoder is used as the low-dimensional input data of the next decoder. Finally, the error is reduced by repeated training layer by layer.

Figure 4 Schematic diagram of stacking noise reduction self encoder (see online version for colours)



4 Utility analysis of music course score clustering based on DAE-K algorithm

4.1 Performance Analysis of DAE-K algorithm

The main idea of DAE-K algorithm is to reduce the dimension of high dimensional data with huge amount of data to improve the clustering effect of K-means algorithm. Therefore, in this study, four datasets, Iris, Wine, Yeast and Statlog image segmentation from UCI standard dataset are selected for performance test comparison and analysis. The parameters of the four datasets are shown in Table 1. The wine dataset records eleven characteristics of red wine, including fixed acidity, volatile acidity, citric acid, residual sugar, chloride, free sulphur dioxide, total sulphur dioxide, density, pH, sulphate, and alcohol. The Iris dataset is a flower classification dataset where each sample contains four features: calyx length, calyx width, petal length, and petal width. The yeast dataset is a multi-label dataset containing 14 types of labels. It has been divided into training and test sets and can be used directly for machine learning, multi-label classification, and more. It can be used in both MATLAB and Python.

Table 1 UCI dataset parameter table

<i>Dataset</i>	<i>Number of samples</i>	<i>Characteristic number</i>	<i>Number of categories</i>
Iris	150	4	3
Wine	178	13	3
Yeast	1,484	8	10
Statlog-Image Segmentation	2,310	19	7

Table 1 shows the number of samples, features and categories of the four datasets iris, wine, yeast and Statlog-Image Segmentation. This study compares DAE-K with AE-K and K-means algorithms, and selects normalised mutual information (NMI), adjusted mutual information (AMI), and Follkes and Mallows Index (FMI) as the performance indicators of the algorithm. The values of these three indicators are all between 0 and 1, and the larger the value, the better the performance of the algorithm. Table 2 shows the clustering index results of the Iris dataset.

Table 2 Iris dataset clustering index results

<i>Dataset</i>	<i>Algorithm</i>	<i>NMI</i>	<i>AMI</i>	<i>FMI</i>
Iris	DAE-K	0.896	0.859	0.834
	AE-K	0.894	0.856	0.821
	K-means	0.893	0.852	0.844

It can be seen from Table 2 that on the Iris dataset, there is no significant difference in the NMI, AMI and FMI indicators of the three algorithms. The index values are 0.893, 0.852 and 0.844. The difference between the first two indicators NMI and AMI is only 0.003 and 0.007, and the DAE-K algorithm is lower than the K-means algorithm on the FMI index by 0.01. The reason for this phenomenon is that the Iris dataset has a relatively small number of samples, features, and categories. So overall, the Iris dataset is not complex, and general clustering algorithms like K-means can achieve good clustering results on such datasets. However, the DAE-K algorithm's dimensionality reduction processing of complex high-dimensional data does not significantly improve the clustering performance of simple datasets. Therefore, further testing of clustering performance is needed on more complex datasets. Table 3 shows the clustering index results of the Wine dataset.

Table 3 Clustering index results of wine dataset

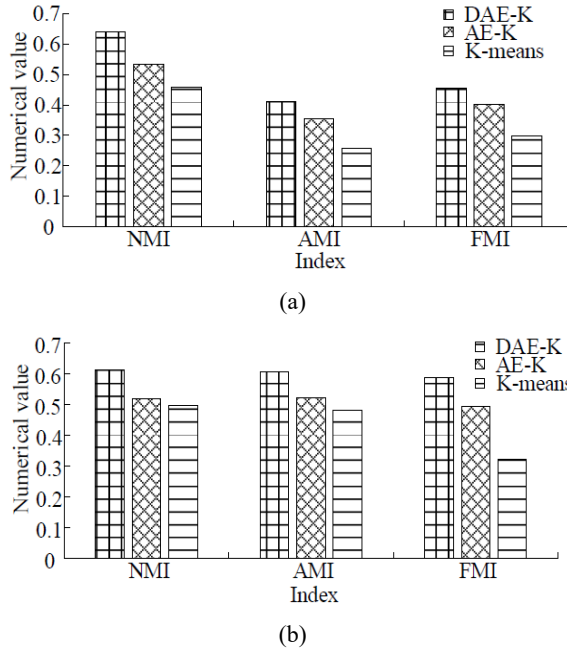
<i>Dataset</i>	<i>Algorithm</i>	<i>NMI</i>	<i>AMI</i>	<i>FMI</i>
Wine	DAE-K	0.900	0.802	0.884
	AE-K	0.897	0.798	0.854
	K-means	0.862	0.762	0.741

From Table 3 that on the wine dataset, the DAE-K algorithm has obtained the optimal index data, the NMI index is 0.900, the AMI index is 0.802, the FMI index is 0.884, and the index value obtained by the K-means algorithm is 0.862, 0.762 and 0.741. Compared with the K-means algorithm, the DAE-K algorithm has a small improvement of 0.038 and 0.04 in NMI and AMI, and a large improvement of 0.143 in the FMI index. Due to the larger sample size of the wine dataset compared to the Iris dataset, the number of features is approximately three times that of the Iris dataset. Therefore, as the complexity of the dataset increases, the DAE-K algorithm can bring better clustering results. Figure 5 shows the indicator histograms of the K-means and DAE-K algorithms on the yeast and Statlog image segmentation datasets.

As can be seen from Figure 5(a), the performance indicators of the three algorithms on the yeast dataset have decreased compared to the iris and wine datasets, because the number of samples in the yeast dataset is much higher than in the Iris and Wine datasets. And the number of categories is also higher, so the indicator has decreased. However, in

comparison, the DAE-K algorithm still obtains good index values and ranks the highest among the three algorithms. Its NMI, AMI and FMI indicators are 0.639, 0.410 and 0.454 respectively, which are higher than the corresponding K-means algorithm. 0.181, 0.152, and 0.157, which are increased by 39.2%, 58.9%, and 52.9%, respectively, when converted into percentages. It can be seen from Figure 5(b) that the DAE-K algorithm has values of 0.614, 0.608 and 0.587 on the NMI, AMI and FMI indicators, while the corresponding index values of the Kmeans algorithm are 0.496, 0.482 and 0.323. The DAE-K algorithm outperforms the K-means algorithm by 23.8%, 26.1%, and 81.7%, respectively. The Statlog Image Segmentation dataset has 2,310 samples, 19 features, and seven categories, making it the most complex dataset among the four datasets. It can be seen that the DAE-K algorithm reduces the complexity of high-dimensional data. It has good effect in dimensional clustering.

Figure 5 Indicator histogram of yeast dataset and Statlog-Image Segmentation, (a) histogram of yeast clustering effect, (b) histogram of Statlog-Image Segmentation clustering effect



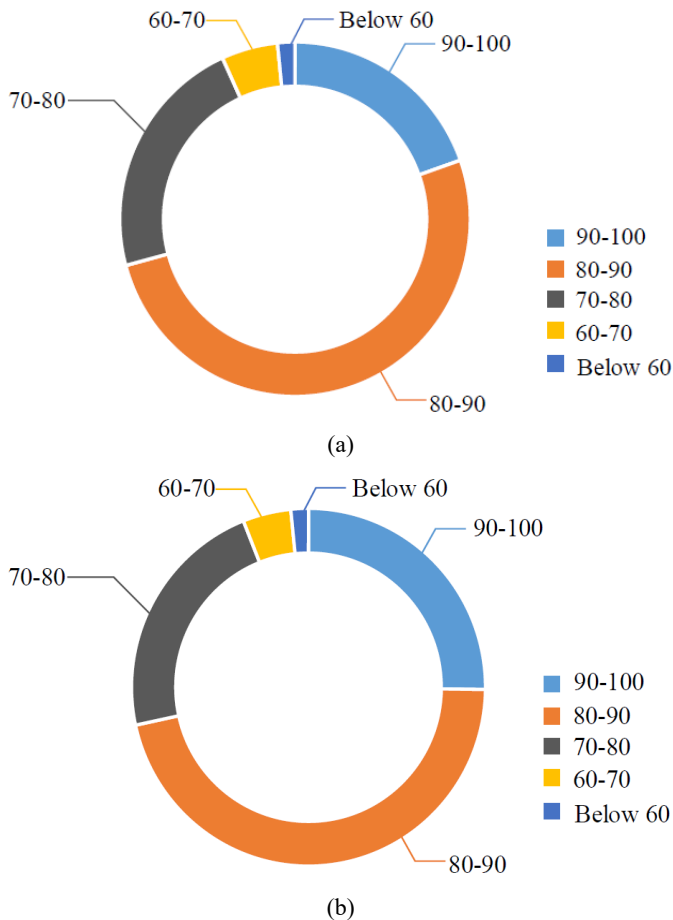
4.2 Analysis of application effect of DAE-K algorithm

In recent years, the country has vigorously strengthened the cultivation of innovative talents, and the requirements for talent cultivation in colleges and universities have become increasingly strict. The analysis of majors and comprehensive qualities of college students by clustering algorithms has become an important tool. This research aims at the training direction of music talents by using data mining of students' past performance and analysing the restrictive factors of the innovation ability of current music majors. The sample selects 500 juniors from a music college in a coastal area, and analyses the score data in two parts: music theory course and practical operation course. The algorithm adopts the DAE-K algorithm proposed in this paper. The music theory course includes

composition analysis and vocal music course, and the practical course includes composition course, improvisation, live performance and listening. Figure 6 shows the distribution of scores in music theory courses.

From Figure 6(a), we can see the distribution of grades in the work analysis course. Most students (over 90%) have a score of 60 or above; students with scores above 80 account for over 50%; a considerable number of students have job analysis scores above 90. This indicates that students have good theoretical performance in job analysis. From Figure 6(b), we can see the distribution of scores of vocal music courses. Among them, most of the students (more than 90%) scored more than 60 points, and the number of students who scored more than 80 points accounted for more than 50%, and the scores were higher than 90 points. The proportion of the number of people is higher than that of the work analysis, which indicates that the level of the students in the vocal music theory class is relatively good. Therefore, in general, the students have a good and excellent development trend in mastering theory courses. Figure 7 reflects the performance results of students in the four courses of practical training.

Figure 6 Score distribution of music theory course, (a) score of composition analysis, (b) vocal music course grades (see online version for colours)



In Figure 7(a), the DAE-K algorithm divides the practice class scores into three categories. Each part represents a score segment, and the distribution of the three categories is relatively concentrated, and the data point sets of different categories have obvious scores. At the same time, there are fewer outliers, so the clustering effect of this algorithm is good. There are differences between different categories, and similar data points have similarity. This also indicates that the clustering results obtained by this algorithm are reasonable for students' practical courses. Figure 7(b) reflects the specific score distribution of the three categories of scores. Among them, the scores of the three categories corresponding to composition are 61.2, 70.5, and 78.5, and the scores of the three categories corresponding to improvisation are 62.5, 72.4, and 79.3. The scores of the three categories corresponding to live performance are 63.2, 76.7, and 82.6, and the scores of the three categories corresponding to listening are 64.6, 71.5, and 80.5. Therefore, compared to music theory courses, the average score of practical courses is lower. This indicates that students have a good theoretical level, but their innovation ability cannot keep up with the pace of theory, and practical training needs to be strengthened. Figure 8 shows the overall score of students in music performance.

Figure 7 Cluster analysis results of practical courses, (a) clustering results of practical courses, (b) scores of practical courses (see online version for colours)

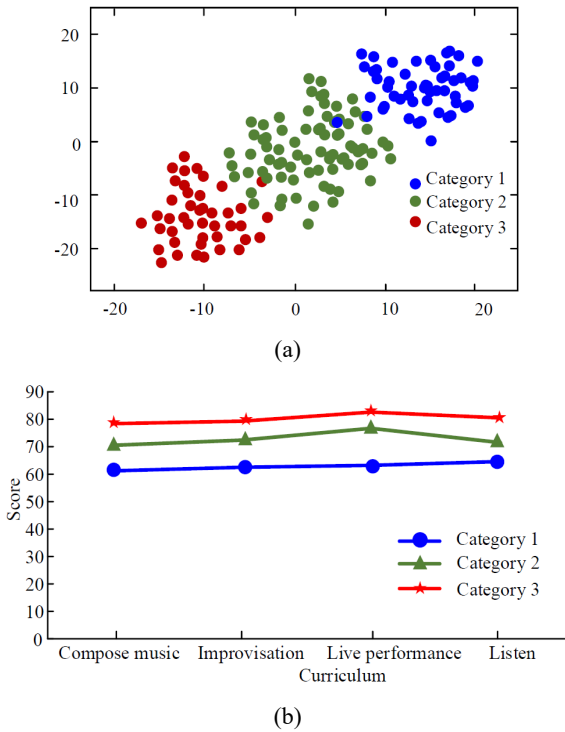
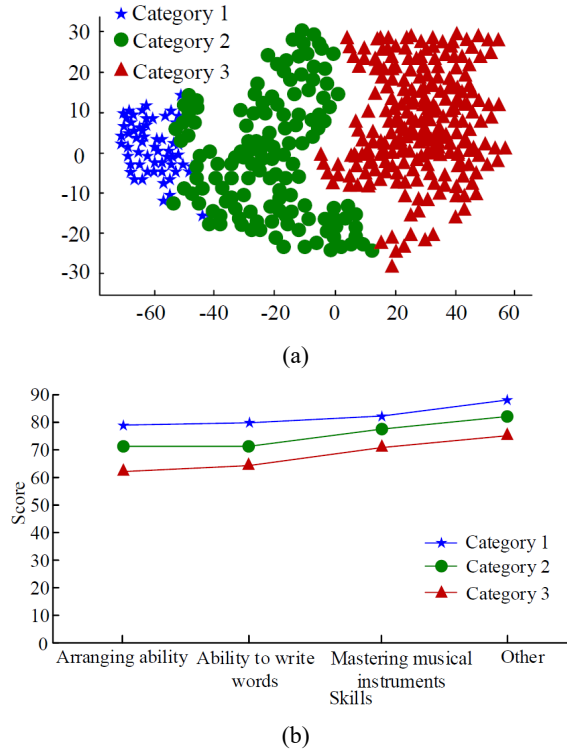


Figure 8 shows that students' overall music scores are grouped into three categories. The first category scores above 80 points in various skills, the second category scores between 70 and 80 points, and the third category scores between 60 and 70 points.

Figure 8 Score of overall music performance, (a) clustering results of practical courses, (b) scores of practical courses (see online version for colours)



5 Conclusions

Innovation ability is very important for talent development, and the demand for innovative talent is always in short supply. At the same time, the level of a country's innovation capability directly affects its sustainable development level. Therefore, it is necessary to study the factors affecting the innovation ability of talents. In this paper, the noise reduction autoencoder is used to reduce the dimension of the music course data, and the K-means algorithm is used to perform cluster analysis on the grade data. The results show that the DAE-K algorithm proposed in this paper is better than the AE-K algorithm and the K-means algorithm. Compared to the K-means algorithm, it improved by 0.038, 0.04, and 0.143, respectively. On the Statlog image segmentation dataset, the corresponding index values of the DAE-K algorithm are 0.496, 0.482 and 0.323, which are 23.8%, 26.1% and 81.7% higher than the K-means algorithm. Therefore, it can be seen that the DAE-K algorithm has a better clustering effect than the K-means algorithm and is suitable for clustering analysis of high-dimensional data. Through the cluster analysis of the performance data of innovative talents in music, it can be seen that the students' performance in music theory class is better than that in music practice class. Among them, in the score distribution of works analysis theory class and vocal music theory class, more than 90% of the students are above the passing grade, more than 50%

of the students are above 80 points, and many students have scores above 90 points, while music practice. The clustering of practice classes found that whether in composition, improvisation, live performance or listening, the average score is more difficult to exceed 80 points, and the scores are more concentrated between 60 and 70 points, so it can indicate that students' performance in music practice. If the ability does not keep up with the knowledge level of music theory, the training and training of students' practical courses should be increased when equating the curriculum training plan. This research only addresses the limitations of K-means combined with noise reduction auto-encoder. In fact, noise reduction auto-encoder is a more general and efficient way to reduce dimensions. You can try to combine noise reduction auto-encoder with other clustering algorithms.

References

- Anggarwati, D., Nurdian, O. and Ali, I. (2021) 'Penerapan Algoritma K-Means Dalam Prediksi Penjualan Karoseri', *Jurnal Data Science & Informatika*, Vol. 1, No. 2, pp.58–62.
- Avik, S., Arnab, A and Ashish Kumar, J. (2022) 'Innovational duality and sustainable development: finding optima amidst socio-ecological policy trade-off in post-COVID-19 era', *Journal of Enterprise Information Management*, Vol. 35, No. 1, pp.295–320.
- Cao, R., Zhang, C., Luo, Y. et al. (2021) 'Comprehensive construction of innovative talent cultivating model in colleges and universities', *Asian Agricultural Research*, Vol. 13, No. 4, p.5.
- Chen, J., Yang, M., Guo, Y. et al. (2021) 'Innovation ability cultivation of graduate students of computer science and technology under the background of science and education integration', *Advances in Applied Sociology*, Vol. 11, No. 12, p.8.
- Cui, M. (2020) 'Introduction to the k-means clustering algorithm based on the elbow method', *Accounting, Auditing and Finance*, Vol. 1, No. 1, pp.5–8.
- Dong, A. and Gao, Q. (2021) 'Framework of the teaching process based on machine learning and innovation ability cultivation in big data background', *Journal of Physics: Conference Series*, Vol. 1992, No. 4, p.042069 (6pp).
- El-Khatib, S.A., Skobtsov, Y.A. and Rodzin, S.I. (2021) 'Comparison of hybrid ACO-k-means algorithm and Grub cut for MRI images segmentation', *Procedia Computer Science*, Vol. 186, No. 11, pp.316–322.
- Feng, Y., Zhao, S. and Liu, H. (2020) 'Analysis of network coverage optimization based on feedback K-means clustering and artificial fish swarm algorithm', *IEEE Access*, Vol. 8, No. 1, pp.42864–42876.
- Ghazal, T.M. (2021) 'Performances of K-means clustering algorithm with different distance metrics', *Intelligent Automation & Soft Computing*, Vol. 30, No. 2, pp.735–742.
- Hou, C., Wu, J., Cao, B. and Jing, F. (2021) 'A deep-learning prediction model for imbalanced time series data forecasting', *Big Data Mining and Analytics*, Vol. 4, No. 4, pp.266–278.
- Lin, J. and Li, H. (2021) 'A hybrid K-means-GRA-SVR model based on feature selection for day-ahead prediction of photovoltaic power generation', *Journal of Computer and Communications*, Vol. 9, No. 11, p.21.
- Ma, J. (2022) 'Russian music education concept and its enlightenment to Chinese music teaching', *Arts Studies and Criticism*, Vol. 3, No. 1, pp.102–105.
- Nithyanandam, G.K. (2020) 'A framework to improve the quality of teaching-learning process – a case study', *Procedia Computer Science*, Vol. 172, pp.92–97.
- Peng, B. and Zheng, X.X. (2020) 'Exploring the training mode of engineering disciplines in Chinese colleges and universities to meet the needs of an innovative society', *International Core Journal of Engineering*, Vol. 6, No. 5, pp.149–151.

- Ramalingam, M. and Thangarajan, R. (2020) 'Mutated k-means algorithm for dynamic clustering to perform effective and intelligent broadcasting in medical surveillance using selective reliable broadcast protocol in VANET', *Computer Communications*, Vol. 150, No. 2, pp.563–568.
- Sang, F. and Xu, Y. (2022) 'Research on the current situation and improvement strategies of aesthetic education in music classrooms in primary and secondary schools in China', *Modern Economics & Management Forum*, Vol. 3, No. 1, pp.14–19.
- Song, C., Li, L., Li, L. et al. (2021) 'Robust K-means algorithm with weighted window for seismic facies analysis', *Journal of Geophysics and Engineering*, Vol. 2021, No. 5, p.5.
- White, R. (2020) 'Authentic learning in senior secondary music pedagogy: an examination of teaching practice in high-achieving school music programmes', *British Journal of Music Education*, Vol. 38, No. 2, pp.1–13.
- Zhang, B. and Sun, W. (2021) 'An innovative study of the cultivation of film and television talents in higher vocational colleges from the perspective of production-teaching integration', *Advances in Applied Sociology*, Vol. 11, No. 11, p.7.
- Zhang, H. and Peng, Q. (2022) 'PSO and K-means-based semantic segmentation toward agricultural products', *Future Generation Computer Systems*, Vol. 126, pp.82–87.
- Zhang, S. and Mao, H. (2021) 'Optimization analysis of tennis players' physical fitness index based on data mining and mobile computing', *Wireless Communications and Mobile Computing*, Vol. 2021, No. 11, pp.1–11.
- Zhang, Y., Lin, Z., Lin, X. et al. (2021) 'A gene module identification algorithm and its applications to identify gene modules and key genes of hepatocellular carcinoma', *Scientific Reports*, Vol. 11, No. 1, p.5517.