

Multimodal Data Mining and Learning Behavior Analysis for Civic Education

Zhenyu Zhan^{1,✉}, Haining Wang¹

¹ School of Marxism, Shandong University, Jinan, Shandong, 250000, China

ABSTRACT

This paper designs a multimodal data mining and learning behavior analysis model for civic education, uses improved clustering and association rule algorithms to analyze the multimodal data obtained from students, mines the basic consumption, learning and life behavior characteristics, and carries out analysis of the students' civic situation in order to take targeted civic education measures. Aiming at the problem that traditional clustering results are greatly affected by the selection of initial clustering centers, Gaussian density function is used to determine the initial clustering centers, and Euclidean distance is replaced by density-sensitive distance to avoid sensitivity to noise and anomalies, which improves the accuracy of the clustering results of students' behaviors. Then we use the FP-Growth association rule algorithm to improve the Apriori construction, recursively and iteratively construct the frequent pattern tree and get the final frequent item set, which improves the efficiency of student behavior data mining. After analyzing the processed student data of a university, it is found that most of the students have low interest in borrowing books, 38.22% of the students borrowed only 2.19 books on average, and the total number of times of book borrowing is only 5.4 times, and the average number of days of single borrowing is 62.3 days, and the school library needs to increase the promotion of students' reading, which can be done through the way of offline book fairs and e-recommendations to improve students' interest in reading books. Reading interest. The study makes a useful exploration for the informatization and intelligentization of ideological education in colleges and universities.

Keywords: fuzzy C-mean clustering, Gaussian density function, FP-Growth, civic education

1. Introduction

With the continuous integration of artificial intelligence, big data and other emerging technologies and the field of education, Chinese education has entered the era of big data-driven intelligence, and optimizing the teaching process with technology has become a new trend in educational research [1-

✉ Corresponding author.

E-mail address: 202220141042@mail.sdu.edu.cn (Z. Zhan).

Received 25 September 2024; Revised 15 November 2024; Accepted 05 March 2025; Published Online 16 April 2025.

DOI: [10.61091/jcmcc127b-025](https://doi.org/10.61091/jcmcc127b-025)

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

2]. Smart classroom is an emerging product of the integration of information technology and education and teaching, and it is also the main scenario for realizing smart education and classroom change [3-4].

The main points of the work of the Ministry of Education clearly pointed out that it is necessary to accelerate the digital transformation and upgrading of education, create digital teaching resources and service platforms, rely on big data technology to realize the mining and processing of the whole process of data in the smart classroom, support the multimodal analysis of the learning situation and teaching intervention decision-making, and promote the reform of the evaluation of smart classroom teaching, and ultimately realize the process optimization of smart teaching [5-8].

Evidentiality is the main feature of the current smart classroom teaching evaluation, and the collection, storage and analysis of students' online and offline learning data based on big data technology can help teachers discover and improve problems in the teaching process in a timely manner, and promote the transformation of educational decision-making from experience-based to data-based [9-11]. Compared with the traditional classroom teaching environment, the smart classroom changes not only the physical teaching environment, but also the subversive change of the teaching process. The smart teaching platform is an important part of the smart classroom teaching environment, which is able to record a large amount of teaching and learning behavioral data, which are characterized by large capacity, multidimensionality, and short time series [12-14]. Mining the teaching behavior data in the smart teaching platform is of great significance in realizing the reform of smart classroom teaching evaluation process [15].

The big data of teaching behavior in the smart classroom teaching platform usually cannot be directly equated with educational information, and it is necessary to rely on data mining technology to explore the hidden potential value behind the data [16-17]. Educational data mining is formed by the intersection of three disciplines: computer science, statistics and education, and usually adopts automated processing algorithms such as classification, clustering, regression and association rule mining to explore the values and patterns contained in large amounts of educational data, which has the unique advantages of supporting the processing of multi-dimensional data, high analytical efficiency, flexible algorithms, and accurate visualization results [18-20]. Analyzing teaching goals based on certain educational big data and tracing the essence of education with the help of flexible algorithms effectively solves the problems of multiple dimensions of student data, large scale, long time span, and difficult to refine the evaluation in the smart learning environment, and avoids the empirical decisions of teachers in the process of educational decision-making [21-22].

Wisdom teaching can, to a certain extent, enrich the teaching form of the Civics and Political Science class, broaden the main channel and main position of ideological and political education, and better utilize its nurturing function [23]. Existing research related to the wisdom learning classroom of the ideological and political class in colleges and universities is mainly based on the macro or meso perspective, and is carried out from the aspects of value orientation, teaching mode, learning mode, learning effect, etc., and there is a lack of research carried out from the perspective of the more microscopic online learning behavior [24-25].

Smart classroom based on data mining technology and artificial intelligence technology is the main position for realizing educational intelligence, and the analysis of students' behavior in the classroom helps teachers to continuously optimize their decision-making and improve the effectiveness of classroom teaching. Jie, W et al. based on the data of students' search logs of self-designed independent learning platform and combined with the decision tree algorithm to analyze the factors affecting students' resource searching behavior and logging in behavior, and used this as the basis for optimizing the independent learning platform and course design [26]. El Aouifi, H et al. analyzed how learners interact with instructional sequences based on educational data mining techniques and found that sequences of learning videos can be used for prediction of students' performance, as well as assisting teachers in understanding students' learning and providing targeted and effective instruction

[27]. Liu, S et al. studied the MOOC online course cloud platform in conjunction with sequential analysis strategy, aiming to provide educators with evidential, analytical, and contextual data results related to students' learning, which can be an effective reference for educators to grasp students' status [28]. Gushchina, O. M et al. synthesized cluster analysis, data analysis and data visualization to design a set of educational data mining techniques for enhancing the effectiveness of e-learning, which can be used to assess students' e-learning behaviors, as well as to understand the dynamics of students' interest in learning and the quality of educational content [29]. Chen, L et al. conceived a data mining algorithm with student behavioral characteristics as the underlying logic and used it in practice to learn that the proposed scheme can efficiently and accurately predict students' learning behaviors and life dynamics, which provides an effective reference for teaching management and teaching reform [30]. Khan, A et al. systematically reviewed the studies related to educational data mining based modeling of student academic performance and pointed out that the present day is effective for predicting grades during the course tenure, while the prediction of grades prior to the commencement of the course is in great need of attention [31]. The above study analyzes the application of data mining technology in the analysis of student behavior, which realizes the prediction of students' performance, and at the same time helps educators to grasp students' behavioral thinking, which provides an important reference for educators to take effective and reasonable educational management measures and educational guidance. Some scholars have also analyzed the feasibility and benefit of data mining technology empowering ideological education. iZhou, W et al. elaborated that data mining technology empowered education field promoted the management of education informatization, and significantly improved the quality and efficiency of teaching, and proposed to scientifically use modern information technology and data mining technology to add support to Civic Education in order to promote the effectiveness of Civic Teaching and Teaching efficiency [32]. Yang, R elaborated on the status quo and potential obstacles of the integration and development of big data and data mining technology with Civic and political education, and pointed out that we should make good use of big data technology to realize the reform and innovation of Civic and political education, which was specifically analyzed in detail from the five dimensions, and made a positive contribution to the innovation and reform of Civic and political education [33]. Scholars all believe that big data as well as data mining technology helps to promote the informatization construction and innovative reform of ideological and political education.

This paper improves the traditional clustering and association rule algorithms to design a multimodal data mining and learning behavior analysis model for Civic Education. Firstly, we analyze the important fields of one-card dataset and achievement dataset, as well as the characteristics of consumption behavior, learning behavior and life behavior. Then the clustering and association rule algorithm is improved. Fuzzy C-mean clustering algorithm with density-sensitive distance is used to optimize the clustering effect by adjusting the calculation of distance. In the traditional Apriori base construction on the generation of frequent itemsets all libraries are abstracted into frequent pattern trees, compression still preserves the itemset association information, optimized as a deformed association rule algorithm. After the design is completed, it is applied on a university student dataset to analyze the students' civic behavior.

2. Method

2.1. Multimodal Data Mining for Civic Education

Data mining refers to a process of searching for the hidden information behind the massive data through some algorithmic models. Data mining is a hot issue under the background of the combination of the new era of Internet technology and artificial intelligence, it is a kind of technology based on the extraction of raw data to understand the basis of analysis, actively search for information behind the data, the effective information will be summarized and organized to gradually find out the law behind

the things. Data mining technology and statistical learning are both sciences that need to seek progress in continuous exploration.

After deciding to adopt the data mining method, it is very necessary to logically take the corresponding steps, understand the specific content of each step in detail, and clearly analyze the research purpose of the experiment. Utilizing data mining techniques to discover problems from cumbersome and complex data and to mine the valuable information behind their problems by building interesting models requires the following steps.

1) Define the problem. First of all the most basic and important thing before starting knowledge discovery is to have a clear perception of the data as well as a basic grasp of the business. We need to establish a clear direction, a comprehensive grasp of the goal, a clear definition. For example, for the problem definition, from the most basic and simple data mining perspective, to determine whether the problem belongs to the regression problem or classification problem.

2) Data mining library. Simply put, in the process of data acquisition, we need to target the specific needs of the project, from the original data source to obtain the relevant data sets we need, the most basic method of database access, online filtering capture, questionnaires and many other forms. At this time, the obtained data set is called the original data set, but due to the data set we obtained there are some inevitable problems, the original data set can not be used directly, data duplication and redundancy, part of the data is incomplete, the significance of missing fields, the outline is not uniform, etc., we will be this type of data set is called noisy, incomplete data. If further operations are carried out directly on the original data, even if the model construction is completed, the final results may have quality problems, and its accuracy cannot be guaranteed. Therefore, before the next step in the operation of the data of the problems that exist in a series of processing, for the improvement of the efficiency of data mining and the accuracy of the final results to provide the necessary protection.

3) Analyze the data. This part is the most critical procedure, we first need to observe the data from various dimensions, try to seek to explore the possible hidden value of the information behind the data, through effective analysis to find out the most influential and reliable factor variables on the results, so as to carry out a new round of definition and discovery.

4) Preparing data. After analyzing the data, a more complete and systematic dataset is often needed, and this step solves the problem in a more perfect way. In this process, variables need to be selected and recorded, even if some new variables need to be created.

5) Modeling. First of all, we need to understand the establishment of the model is a process that needs to be repeated, for different problems need to choose a different model for fitting, the choice of the model can not be done overnight, but should be gradual, slowly explore, until we find the most appropriate and most valuable model, to get a good prediction of the fitting results.

6) Evaluate the model. After establishing a good model, how to judge whether the model we chose before is optimal, then we need to evaluate the model to achieve the purpose. Evaluating the model can make the experimental argumentation results more convincing and reliable.

7) Implementation. After completing the above steps, the next step is to practice the effective results obtained from data mining. The first one is to give some feasible suggestions to the relevant personnel; the other is that if the model conditions are good, you can further consider applying this model to other problems, which not only makes the model use more rich and mature, but also makes the understanding of different problems more profound and clear.

In summary, the multimodal data mining process is shown in Figure 1.

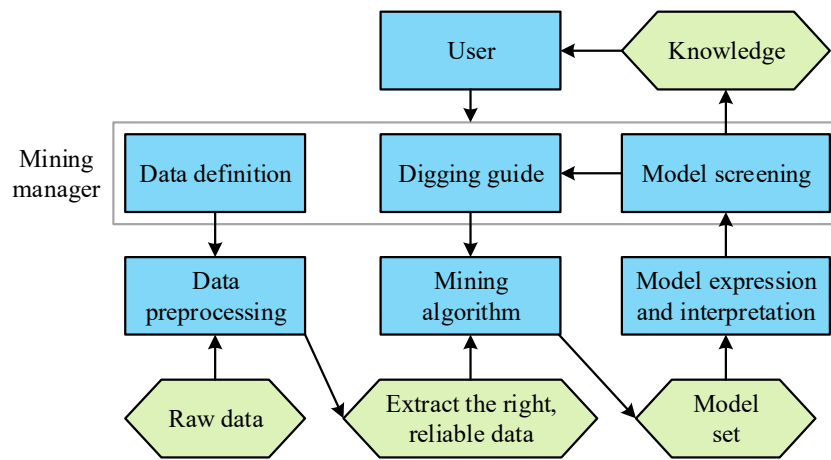


Fig. 1. Schematic diagram of data mining and extraction process

2.2. Data pre-processing

This section elaborates on the dataset required for the research content of this paper. The behavioral characteristics of students mainly include three categories: consumption behavioral characteristics, learning behavioral characteristics and life behavioral characteristics. In the acquired data, starting from students' basic information, consumption data and performance data, the most valuable data in students' campus activities are selected for feature extraction through categorization and summarization.

2.2.1. Student consumption dataset. There are four sources of overall data obtained in this paper, which are "One Card" consumption data, student basic information data, access control data, and achievement data, of which the main sources are "One Card" database and achievement database.

The One Card database has all the consumption records of undergraduates, graduate students, faculty and staff, as well as foreign temporary employees, but in the study of accurate student financial aid in this paper, it is necessary to judge the economic situation.

2.2.2. Student Achievement Dataset. The student achievement data comes from the achievement database, which contains the course achievement data of the enrolled students between the last semester of 2021 and the next semester of 2022, and the data contains 19 colleges, 94 majors, 1,268 courses, and 5,022 students, with a total of 287,737 pieces of data, and the data fields are student number (xh), class (BJ), major name (zymc), course grades (kccj), academic year (xn) and so on.

Student behavioral characteristics are extracted from the acquired data, which are consumption behavioral characteristics, learning behavioral characteristics and life behavioral characteristics.

2.2.3. Characteristics of consumer behavior. The extraction of consumption behavior characteristics is mainly based on the data from the campus "one card" database, which mainly records the consumption of the campus cafeteria, super while, water, etc. In this paper, we distinguish consumption into cafeteria and other consumption. According to the needs, this paper distinguishes consumption into cafeteria consumption and other consumption. Extract the chronological characteristics of student consumption and calculate the number of consumption of school students. Analyze the characteristics of consumption and extract the basic consumption characteristics, which mainly include total number of consumption, total amount of consumption, average value of each consumption, average value of monthly consumption, number of weekly consumption and number of monthly consumption as the basic consumption behavior characteristics.

2.2.4. Characteristics of learning behavior. In this paper, the degree of stability of academic performance W is extracted from the performance database and the formula is shown in (1):

$$W = \tanh \frac{\sum_{i=1}^n (SC_{sd_i} - \overline{SC_{sd}}) \times N_{m_i}}{\sum_{i=1}^n N_{m_i}} \quad (1)$$

where: n is the number of performance grade interval classifications; N_{m_i} is the number of times the student achieved the i th category grade; $\overline{SC_{sd}}$ represents the standard deviation of all performance grade classification interval grades, and SC_{sd_i} represents the standard deviation of the i th category grade.

In addition, for the student's academic performance for the academic year, the semester was used as the calculation interval to define the semester academic performance excellence G , which is calculated as shown in formula (2)

$$G = \frac{(\sum_{i=1}^n (S_i - R_i)) \times F_1}{\sum_{i=1}^n F_i \sum_{i=1}^n S_i} \quad (2)$$

where: F_i denotes the number of credits in course i ; n denotes the total number of courses taken in the semester; R_i denotes the grade rank of course i , and S_i denotes the class size.

2.2.5. Characteristics of Life Behavior. In this paper, the characteristics of life behaviors extracted from the One Card database are regularity of meals, regularity of waking up in the morning, regularity of bedtime, and regularity of eating places. Studies have shown that regular living habits will have an impact on students' academic performance.

1) Dining Regularity: It means that in the breakfast, lunch and dinner time, the meals are taken steadily and on time, and the regularity coefficient is recorded as λ , and the expression is shown in (3):

$$\lambda = \frac{M}{3 \times D} \quad (3)$$

Where: M is the number of times the student ate during breakfast, lunch, and dinner periods; and D indicates the number of days the student was in school.

2) Early rising regularity: calculate the probability of the student's early rising occurring during the school period. In this paper, we can not know the specific time of the student's early rise, the breakfast meal time to roughly estimate whether the student early rise, the degree of regularity of early rise is recorded as z , the mathematical expression as shown in (4):

$$Z = \frac{N}{D} \quad (4)$$

where: N denotes the number of breakfast meals; D denotes the number of days in the time period taken.

3) Bedtime regularity: the projection of bedtime in this paper is calculated by the last hot water use time at night, but it is only used as a rough estimate because not every student uses hot water every night and does not rest after the last use of hot water. The degree of bedtime regularity is denoted as Q , and the mathematical expression is shown in (5):

$$Q = \frac{N}{M - T_{\min}} (T_{\max} - T_{\min}) \quad (5)$$

where: N indicates the number of early bedtime; M indicates the average time to sleep; $T_{\min}$ indicates the earliest bedtime, this paper takes 18; $T_{\max}$ indicates the latest bedtime, this paper takes the next day at 3:00 a.m., the value is converted to 27.

4) Regularity of dining places: the “one card” data obtained by the school, the campus normal supply of dining cafeteria has 7, distributed in all parts of the campus, but the same professional grade study area is more concentrated, so the frequency of dining places can be a certain degree of reaction to the scope of the student's learning activities. The frequency of dining places is recorded as Fr , and the mathematical expression is shown in (6):

$$Fr = \sqrt{\frac{1}{n} \sum_{i=1}^n (DR_i - \overline{DR})^2} \quad (6)$$

where: n denotes the number of cafeterias that students have consumed, which in this paper is a positive integer less than or equal to 7; DR_i denotes the logarithm of the number of times consumed in the cafeteria; and $\overline{DR}$ denotes the average of the logarithm of the number of times consumed in the cafeteria.

2.3. Fuzzy C-mean clustering algorithm improvement

2.3.1. Gaussian Density Function to Determine Initial Cluster Centers. In traditional FCM, the initial clustering centers are usually chosen randomly, and this practice may result in clustering results that are highly influenced by the choice of initial clustering centers [34]. In order to improve this situation, a Gaussian density function can be used to determine the initial clustering center. Specifically, the following steps can be followed:

1) Calculate the density of each data point. A Gaussian density function can be used to calculate the density of each data point, given a student dataset $S = \{x_1, x_2, \dots, x_n\}$, where the density of each object x_i is denoted as $density(x_i)$ is the average of the sum of Gaussian distances of the data points around x_i , which is used as a measure of the localized density at which x_i is located, and the Gaussian function is defined as shown in Eq. (7).

$$density(x_i) = \sum_{j=1}^n f_{Guass}(x_i, x_j) \quad (7)$$

$$f_{Guass}(x_i, x_j) = e^{-\frac{d(x_i, x_j)^2}{2\sigma^2}} \quad (8)$$

where σ is the density parameter.

2) Select the data point with the largest density as the initial clustering center. Sort all the data points according to the density and select the data point with the largest density as the initial clustering center [35]. In order to avoid the calculated density maximum points are all in one cluster, choose the appropriate region length to divide, the formula is as follows:

$$d = \frac{D_{\max}}{K} \quad (9)$$

where the number of clustering categories K , $D_{\max}$ is the maximum value of the distance between two points among all data points. The algorithm for determining the initial cluster centers using Gaussian density function is as follows:

Input: data set S , number of clusters K

Output: set of initial clustering centers $C = \{c_1, c_2, \dots, c_k\}$

Step 1: Calculate and determine the maximum distance between data points in the data space $D_{\max}$, and the constraint distance d ;

Step 2: Find the point c_1 with the highest density among all data points.

Step 3: Find the 2nd densest point c_2 with distance c_i greater than d ;

Step 4: Keep finding cluster centers of other initial clusters as per step 3 $c_3, \dots, c_k$;

Step 5: Return the set of cluster center vectors $C = \{c_1, c_2, \dots, c_k\}$.

2.3.2. Density-sensitive distances. By replacing the Euclidean distance with the density-sensitive distance, the similarity and density distribution between data points can be better reflected, thus improving the accuracy of the clustering results. At the same time, this method can also avoid the defects of Euclidean distance in traditional FCM, such as the problem of sensitivity to noise and outliers. Given two data points i and j , their density-sensitive distance D_{ij}^ρ can be calculated by the following formula:

$$D_{ij} = \min_{p \in P_{ij}} \sum_{k=1}^{|p|} (e^{\rho d(p_k, p_{k+1})} - 1) \quad (10)$$

$$D_{ij}^\rho = \frac{1}{\rho^2} \ln(1 + D_{ij})^2 \quad (11)$$

P_{ij} is the set of all paths connecting points i and j , where path p is a sequence of points consisting of a series of neighboring sample points, the length of path p is the sum of the distances between neighboring points on the path, and $d(p_k, p_{k+1})$ denotes the Euclidean distance between points p_k and p_{k+1} ; $\rho > 0$ is a parameter controlling the degree of sensitivity of the density, and a smoothing and scaling adjustment of the distances is obtained by the introduction of ρ , which makes the dense region within the the distance metric between the points in the dense region to be smaller.

The modified FCM algorithm uses the density-sensitive distance as the distance metric, and its objective function is shown in Equation (12):

$$J = \sum_{i=1}^n \sum_{j=1}^k u_{ij} D_{ij}^\rho \quad (12)$$

where, D_{ij}^ρ is the density sensitive distance between the sample x_i and the clustering center c_j and u_{ij} is the affiliation of the sample point x_i in the j th grouping, the formula for obtaining the new affiliation matrix by minimizing the objective function while satisfying the constraints is shown in the following equation:

$$\sum_{j=1}^k u_{ij} = 1 \quad (13)$$

$$u_{ij} = \frac{1}{\sum_{c=1}^k \left(\frac{D_{ij}^\rho}{D_{ic}^\rho} \right)^{\frac{2}{m-1}}} \quad (14)$$

By obtaining D_{ij}^ρ through the density-sensitive distance, the objective function J can be expanded as:

$$J = \sum_{i=1}^n \sum_{j=1}^k u_{ij} D_{ij}^{\rho} = u_{j1} D_{j1}^{\rho} + u_{j2} D_{j2}^{\rho} + \cdots + u_{jn} D_{jn}^{\rho} \quad (15)$$

The density-sensitive distance matrix D^{ρ} between any two sample points is known:

$$A = \begin{bmatrix} D_{11}^{\rho} & D_{12}^{\rho} & \cdots & D_{1n}^{\rho} \\ D_{21}^{\rho} & D_{22}^{\rho} & \cdots & D_{2n}^{\rho} \\ \vdots & \vdots & & \vdots \\ D_{n1}^{\rho} & D_{n2}^{\rho} & \cdots & D_{nn}^{\rho} \end{bmatrix} \times \begin{bmatrix} u_{11} & u_{21} & \cdots & u_{k1} \\ u_{12} & u_{22} & \cdots & u_{k2} \\ \vdots & \vdots & & \vdots \\ u_{1n} & u_{2n} & \cdots & u_{kn} \end{bmatrix} = \begin{bmatrix} a_{11} & a_{12} & \cdots & a_{1k} \\ a_{21} & a_{22} & \cdots & a_{2k} \\ \vdots & \vdots & & \vdots \\ a_{n1} & a_{n2} & \cdots & a_{nk} \end{bmatrix} \quad (16)$$

Then the clustering center c_j of class j is:

$$c_j = \{x_i \mid i = \arg \min A_{cj}\}, j = 1, 2, \dots, k \quad (17)$$

The specific steps of the overall algorithm are as follows:

Algorithm 3: Fuzzy C -mean clustering algorithm based on density-sensitive distance measure

Input: dataset D , number of clusters c , convergence threshold ε Output: set of cluster centers C , affiliation matrix U Step 1: Initialize the cluster centers using Gaussian density based function.

Step 2: For each data point, calculate the density sensitive distance D_{ij}^{ρ} and initialize the affiliation matrix U based on the density sensitive distance.

Step 3: Update the cluster center C . Step 4: Update the cluster center U . Step 5: Repeat steps 3 to 4 and stop the iteration when the affiliation matrix U no longer changes or reaches the predefined convergence threshold.

2.3.3. Determination of the optimal number of clusters. The number of clusters is an important parameter in clustering algorithms, which directly affects the quality of clustering results. There are many ways to determine the optimal number of clusters, such as the elbow rule and contour coefficient method. But alone then use there is always the optimal search of a stable, computational complexity and other issues, so this paper decides to m to determine the optimal number of clusters using a combination of the two Wan style method.

The contour coefficient (SC) is an index used to evaluate the clustering results, and its value is taken between -1 and 1, the larger the value indicates that the clustering results are more reasonable. The contour coefficient takes into account the tightness within clusters and the separation between clusters, and can be used to compare the performance of different clustering algorithms on the same dataset and determine the optimal number of clusters.

The method of calculating the profile coefficients is relatively simple, only need to calculate its $a(i)$ and $b(i)$ for each sample, and then substitute into the formula to get the corresponding coefficients. In the actual calculation, the iterative method is usually used for calculation to improve efficiency. Specifically, the following steps can be used for calculation:

- 1) For each sample i , calculate its average distance $a(i)$ from other samples in the same cluster.
- 2) For each sample i , calculate its average distance from all samples in the nearest other clusters $b(i)$.
- 3) For each sample i , calculate its profile factor $s(i)$.
- 4) Calculate the average contour coefficient of all samples as an evaluation index of the clustering effect.

$$s(i) = \frac{b(i) - a(i)}{\max\{a(i), b(i)\}} \quad s(i) = \begin{cases} 1 - \frac{a(i)}{b(i)}, & a(i) < b(i) \\ 0, & a(i) = b(i) \\ \frac{b(i)}{a(i)} - 1, & a(i) > b(i) \end{cases} \quad (18)$$

$a(i)$ indicates the closeness of sample i to other samples in the same cluster, and $b(i)$ indicates the separation between sample i and other clusters. The value of the contour coefficient ranges between $[-1, 1]$, the larger the value, the better, when the coefficient is close to 1, it indicates that the clustering effect is good; when the coefficient is close to 0, it indicates that the clustering effect is average; when the coefficient is close to -1, it indicates that the clustering effect is poor.

The elbow rule is a method used to determine the optimal number of clusters in a clustering algorithm. When the dataset is divided into different clusters, the sum of squares within clusters (SSE) decreases as the number of clusters increases. Thus, when the number of clusters increases to a certain point, the rate of decrease in SSE decreases dramatically, creating an “elbow” point that corresponds to the optimal number of clusters.

The use of the elbow rule is to determine the optimal number of clusters in a clustering problem. Define the SSE formula as follows:

$$SSE = \sum_{i=1}^k \sum_{j \in C_i} |j - m_i|^2 \quad (19)$$

where m_i is the clustering center of the i th class and j is the data points belonging to the i th class, not all points do the distance operation with all clustering centers, otherwise the SSE will not converge to 0 with the increase of the number of clusters.

2.4. Association Rule Algorithm and Improvements

2.4.1. Introduction to correlation coefficients. The correlation coefficient generally refers to a coefficient that measures the linear correlation between the factors of all the variables, and was initially devised by the statistician Cole Pearson as a statistical criterion, usually denoted by r . There is more than one definition of correlation coefficient, often depending on the category of the object of study, and most of them use Pearson's correlation coefficient, which can also be called Pearson's product-moment coefficient. The correlation coefficient is introduced as a quantitative measure of the closeness of correlation between variables.

Correlation coefficients are generally calculated using the product difference, which first calculates the deviation between the mean of the variable and itself, and then multiplies the deviations to quantitatively measure the degree of correlation between the objects of study, with a focus on a single correlation coefficient that is linear. By the diverse attributes of these research objects, the name of the standard way of statistics is not the same. The correlation between the assessed variables is usually denoted as $r(X, Y)$, which is calculated as shown in equation (20):

$$r(X, Y) = \frac{\sum_{i=1}^n (x_i - \bar{X})(y_i - \bar{Y})}{n\sigma_X\sigma_Y} \quad (20)$$

In the above equation, n denotes the number of tuple pairs of variable X, Y , x_i and y_i denote the values on variable X, Y on the i th set of tuples, $\bar{X}$ and $\bar{Y}$ denote the mean of variable X, Y , and σ_X and σ_Y are the standard deviation of variable X, Y , which are defined as follows:

$$\sigma_X = \sqrt{\frac{1}{N} \sum_{i=1}^n (x_i - \bar{X})^2} \quad (21)$$

$$\sigma_Y = \sqrt{\frac{1}{N} \sum_{i=1}^n (y_i - \bar{Y})^2} \quad (22)$$

In the above equation, N is equal to $n-1$, n is the number of tuple pairs of variable X, Y , x_i and y_i denote the values on variable X, Y on the i th tuple, and $\bar{X}$ and $\bar{Y}$ denote the mean values of variable X, Y , respectively.

Person correlation coefficient $r(X, Y)$ generally ranges from -1 to 1, i.e., $-1 < r(X, Y) < 1$. When $r(X, Y) > 0$, variable X, Y is positively correlated. That is, the value of variable X increases as the value of variable Y increases, and the larger the value of $r(X, Y)$, the closer the degree of X, Y positive correlation; when $r(X, Y) < 0$, variable X, Y is negatively correlated, that is, the value of variable X decreases as the value of variable Y increases. When $r(X, Y) = 0$, X, Y is uncorrelated, i.e., the change between the two does not imply Y a change with X .

2.4.2. Distributed FP-Growth Association Rules. FP-Growth algorithm is an association analysis algorithm that mainly uses the following kind of strategy: to abstract all the libraries that generate frequent itemsets into frequent pattern trees (FP trees), while compressed yet still preserving the itemset association information [36]. It is mainly improved on the traditional Apriori algorithm and optimized as a deformed association rule algorithm. There is a feature in the Apriori algorithm: all non-empty subsets of the frequent itemsets must also be frequent. However, the performance of Apriori algorithm is usually poor when it is used to mine long frequent patterns, so FP-Growth algorithm is proposed. In this algorithm, the frequent pattern tree is used, which is a kind of prefix tree, but it is very uncommon, mainly consists of two parts, the frequent item header table and the item prefix tree. FP-Growth association rule algorithm utilizes the above construction, which improves the efficiency of mining, and it only talks about two database scans, recursively and iteratively constructs the frequent pattern tree and obtains the final frequent itemset, which improves the efficiency of data mining in terms of time, space complexity, and the efficiency of data mining. The main algorithm idea used in FP-Growth algorithm is to iteratively construct the FP tree and project it, the main steps of the algorithm are as follows:

- 1) Construct the conditional projection database and project the FP tree for each frequent item.
- 2) Construct the newly constructed FP tree iteratively until it is concluded that the new FP tree is empty, or there is only a certain path.
- 3) At the end of the second step of constructing the FP tree, the prefix of this tree is called the frequent pattern; with one and only one path, the cooperation that all seem to occur can be derived, and then the frequent pattern is derived by connecting with the prefix of this tree.

3. Results and Discussion

3.1. Cluster Analysis of Students' Civic Behavior

After these data were cleaned, integrated, transformed and selected through data preprocessing techniques, the three classes of student behavior indicator data were clustered and analyzed separately according to the fuzzy C-mean clustering algorithm designed in the previous section. As the algorithm is given the number of class clusters for data clustering, its initial cluster center selection is very unstable. Therefore, in this paper, the data are clustered several times on the basis of determining the number of class clusters, and the clustering results are reduced to two dimensions using the t-SNE dimensionality reduction algorithm for visualization.

3.1.1. Cluster analysis of student consumption behavior. The cluster centers and cluster labels for each category of clusters in the obtained clustering results are shown in Table 1. As can be seen from Fig. 2, when the density threshold is 17606, the number of cluster centers obtained according to the initial cluster center finding algorithm is 6. Figure 3 shows the effect of different densities tv and the number of class clusters K on the evaluation index $eval$ of clustering effect, at this time, the improved FCM clustering algorithm has the optimal evaluation index and the best clustering effect. Therefore, the student consumption behavior data is divided into 6 classes.

Table 1. Student consumer behavior data clustering results

Cluster number	Student ratio	Monthly consumption		Weekly	Daily	Single consumption average	Month peak	Comparison result
		Amount	Frequency	Amount	Amount			
1	18.90%	67.29	34.88	15.61	2	0.16	136.56	0
2	22.08%	378.09	124.06	93.63	12.72	1.47	537.7	111101
3	3.06%	98.93	15.12	22.9	3.41	4.62	466.31	10
4	26.39%	206.47	74.76	50.53	7.2	2.09	305.41	0
5	19.48%	91.42	24.54	22.19	2.42	2.2	207.2	10
6	10.09%	638.3	198.87	159.29	22.48	2.27	1564.4	111101
Mean		247.62	77.8	61.28	8.59	2.23	535.66	

From the figure, it can be seen that the students in the first cluster belong to the group of students with exceptionally low consumption, and observing the cluster heart of this type of cluster, it can be found that all the indicators of the cluster heart are much lower than the average value. In addition, the average daily consumption amount, average consumption amount and peak consumption value in a month are all much lower than the other clusters. The above shows that the students in this cluster not only spend very low amounts of money on campus, but are also very likely to spend a lot of money off-campus.

Students in the second cluster are moderate spenders and, with the exception of the sub-average spending amount indicator, are slightly higher than the average of all indicators for the overall student population. Observation of the average number of times per month and the average amount of sub-consumption shows that the students in this cluster consume steadily and frequently and that the students in this type of cluster have an average level of consumption.

Observing the center of the third cluster, it can be seen that the average monthly consumption is lower than the other clusters, and the average monthly consumption is higher than the other clusters, but the average monthly consumption of the center of this cluster is only 98.0 yuan. The above situation shows that this cluster belongs to the student group of extra-low consumption, and the consumption fluctuates greatly and is very irregular. Therefore, it cannot be ruled out that students in this cluster often spend money outside the home.

The fourth cluster of students compared to the overall students of the average value of each indicator, although the average monthly consumption level of students in this cluster is lower, but its average monthly consumption times are higher, which can be inferred that the students in this cluster consume more at school, and the average amount of sub-consumption is lower, which reflects that the students in this cluster live a thrifty life from the side. Therefore, this cluster belongs to the student group with low consumption and more stable consumption.

The fifth cluster is similar to the first and third clusters in that the average monthly consumption amount of the cluster's hearts does not exceed \$100, and the average daily consumption amount and the average sub-consumption amount of the fifth cluster's hearts are almost the same. In other words, the students in this cluster may only spend money at school once a day, and the average number of times they spend money at school in a month is also small, only about 24 times. In other words, students in this cluster may only spend once a day at school, and their average monthly consumption

is small, only about 24 times. The above shows that students in this cluster belong to the extra-low consumption group, and their consumption is small and stable.

Observation of the sixth cluster heart shows that the students in this cluster are much higher than the average value of all the indicators of the overall students, except for the average amount of consumption. Therefore, the consumption level of the students in this cluster is much higher than the average consumption level of the overall students, and their living conditions are superior, and they belong to the group of high-consumption students with stable and frequent consumption.

In addition, it can be found that about 32.17% of the students in the school are above the medium level of spending, while the majority of the students are below the low level of spending. About 22.08% of the students have an average monthly spending amount that fluctuates around 378.09 yuan; about 10.09% of the students have an average monthly spending amount that fluctuates around 638.3 yuan. Thus, the school has a majority of students who are below average spenders and have a low average number of monthly purchases. Only a small number of students have the habit of spending money at school regularly. The student management department of the university needs to appeal to students to develop the habit of spending money at school, and it can increase the motivation of students to eat at school by discounting the price of food in the cafeteria and improving the quality of food in the cafeteria.

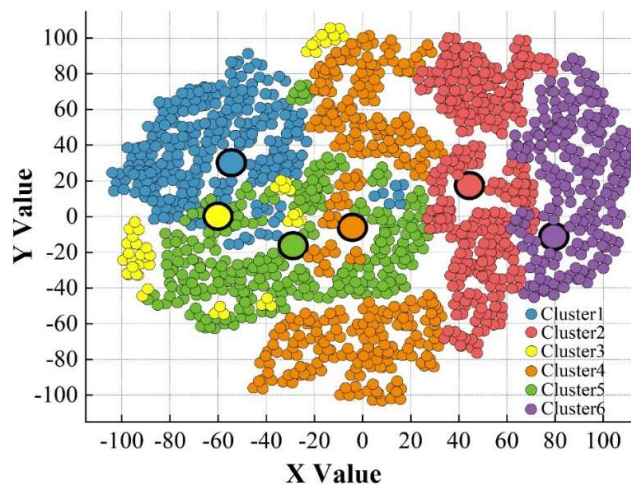


Fig. 2. Student consumption behavior data cluster division

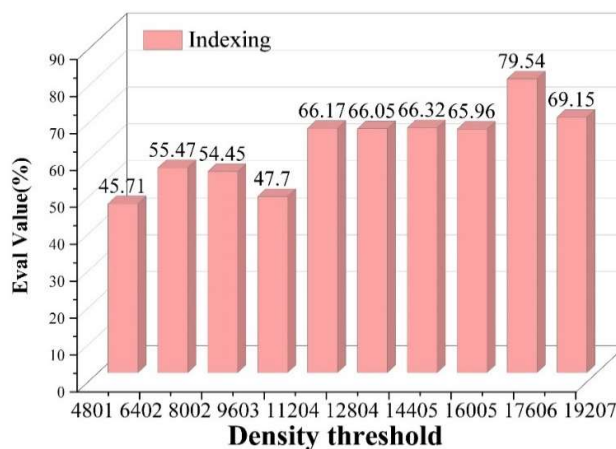


Fig. 3. Evaluation results of student consumption behavior data

Comparison with the evaluation indexes of K-Means algorithm and K-Means++ algorithm is shown in Table 2, the evaluation indexes of K-Means and K-Means++ clustering algorithm under the same

conditions eval are lower than the improved FCM clustering algorithm in this paper. In summary, the improved algorithm of this paper is more advantageous for clustering student consumption behavior data under the same conditions, and the method improves the applicability of K-Means and KMeans++ algorithms for clustering student consumption behavior data.

Table 2. Clustering algorithm evaluation comparison

Algorithm	eval indexing (%)
K-Means	63.11
K-means + +	67.89
This model	77.25

3.1.2. Cluster Analysis of Students' Book Borrowing Behavior. Table 3 shows the clustering results of borrowing behavior. Cluster labels 1~4 are low borrowing, more regular borrowing habits, the most borrowing, more regular borrowing habits, low borrowing, and high single borrowing time, the least amount of borrowing, and unstable borrowing cycle.

Figure 4 shows the impact of different density thresholds and the number of class clusters K on the clustering effect evaluation index eval. The smallest proportion of readers is class cluster 4, accounting for 5.01% of the total number of readers, observing the total number of books borrowed by the cluster center, the total number of times of book lending and the average number of days of single borrowing indicators can be found that the average number of days of single borrowing of the cluster readers is as high as 146.04 days, but the total number of book lending and the total number of times of book lending is the lowest of all the classes of clusters. Thus, the readers of this cluster have the least amount of borrowing and the borrowing cycle is unstable.

Class Cluster 3 accounts for 38.22% of the total number of readers, the average total number of books borrowed by readers in this cluster is only 2.19, the total number of books borrowed is only 5.4 times, and the average number of days of single borrowing is 62.3 days, and it can be inferred from the comprehensive indicators of students' book lending behaviors that the readers in this cluster have a low amount of lending and a high single borrowing time.

Cluster 2 accounts for 8.62% of the total number of readers. Observing the indicators in Cluster 2, it can be found that the total number of books borrowed, the total number of times of book borrowing, and the total number of times of book borrowing are much higher than the other clusters, and the average number of days of single-time borrowing is 40.76 days. Comprehensively, it can be deduced from the above indicators that the readers in this cluster often borrow books, and their borrowing habits are more regular.

Most of the students in this school have a low interest in borrowing books, while only a small number of students have a high interest in borrowing books and a high utilization rate of books. The library of this school needs to increase the promotion of students' reading, which can be done through the offline book fairs and e-recommendations to increase students' interest in reading books.

Table 3. The result of clustering of the behavior index system of raw books

Cluster number	Student ratio	Book borrowing	Total number of books borrowed	Average number of days	Total number of books	Comparison result
1	48.15%	3.39	4.49	22.53	33.24	0
2	8.62%	14.14	34.19	40.76	225.98	1101
3	38.22%	2.19	5.4	62.3	42.81	0
4	5.01%	2	1.91	146.04	33.42	10
Mean		5.23	12.74	68.14	85.28	

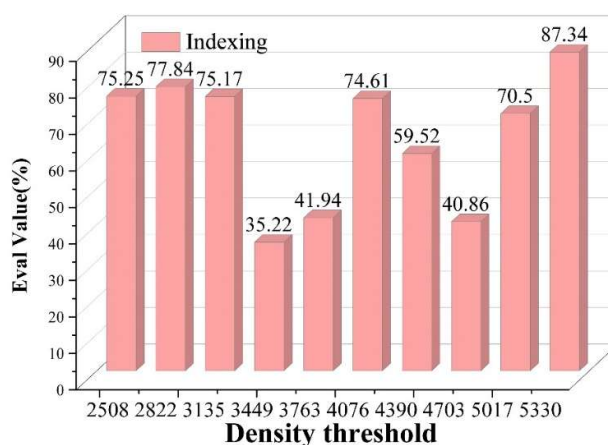


Fig. 4. The results of the evaluation of the student's book borrowing behavior data

The comparison of evaluation indexes with K-Means algorithm and K-Means++ algorithm is shown in Table 4. This paper's algorithm is more advantageous than K-Means and K-Means++ algorithms for clustering student book borrowing behavior data under the same conditions, and the method improves the applicability of K-Means and K-Means++ algorithms for clustering student book borrowing behavior data.

Table 4. Clustering algorithm evaluation comparison

Algorithm	eval indexing (%)
K-Means	73.25
K-means + +	76.34
This model	87.08

3.1.3. Cluster Analysis of Student Exercise Behavior. Clustering was performed using the improved FCM clustering algorithm and the clustering results were obtained as shown in Table 5. Figure 5 shows the effect of different density thresholds and the number of class clusters K on the clustering effect evaluation index eval.

The sports data of school students are divided into five classes, of which the largest proportion of students is class cluster 5, accounting for 34.42% of the total number of sports students, and the cluster heart of this class is larger than the average of the overall students' sports behavior indexes in the three indexes, and a comprehensive observation of the cluster heart of the indexes can be found that the students in this class cluster belong to the group of students who often exercise and often take attendance, and the sports habits are better.

The smallest proportion of students is class cluster 2, which has only one student, accounting for 9.96% of the total number of exercise students. Observing the indicators of attendance, running kilometers, and comprehensive performance of the cluster heart, it can be found that these three indicators are the lowest among all the class clusters. Therefore, it can be inferred that the students of this cluster rarely have the habit of exercising.

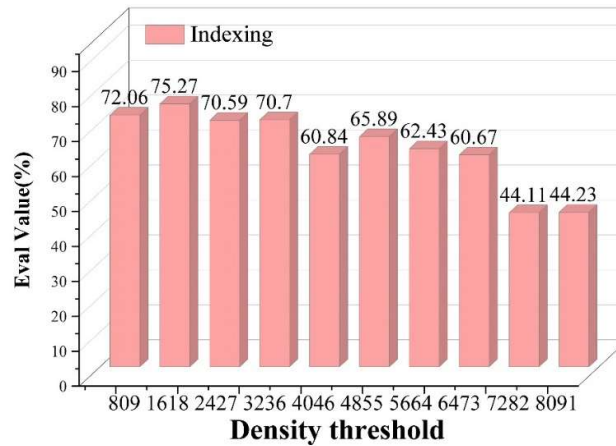
Class Cluster 1 accounted for 19.17% of the total number of athletic students, and the average number of attendance of cluster hearts in this cluster was only 0.17, but their running kilometers and overall performance indicators were the highest among all class clusters. Therefore, the students in this category cluster belong to the group of students who exercise frequently although they rarely take attendance, and student managers should urge the students in this category cluster to participate in attendance regularly.

Cluster 3 accounted for 16.82% of the total number of students in sports. Observing the indicators in cluster 3, it can be found that the heart of this cluster is in the middle of all the clusters in terms of the number of kilometers of running and the overall performance indicators, but the number of times of their attendance is on the low side. Therefore, the students in this cluster belong to the group of moderately athletic students, and the student administrators also need to urge the students in this cluster to participate in the attendance regularly.

Cluster 4, which accounts for 19.63% of the total number of athletic students, has low values for kilometers run and overall performance, but when observing its attendance indicator, it has the highest value of all the clusters in this category. Therefore, it can be inferred that the students in this cluster belong to the group of students who exercise occasionally, and there may be a situation in which they participate in attendance but do not participate in exercise, and student administrators should pay attention to the exercise situation of the students in this cluster.

Table 5. Student movement behavior index system clustering results

Cluster number	Student ratio	Attendance frequency	Running mileage	Comprehensive achievement	Comparison result
1	19.17%	0.17	117.43	76.21	11
2	9.96%	0.97	10.91	8.44	0
3	16.82%	0.66	70.52	48.46	11
4	19.63%	12.29	35.98	32.59	100
5	34.42%	11.69	72.9	58.37	111
Mean		5.26	61.69	44.64	

**Fig. 5.** The results of the evaluation of student sports behavior data

The evaluation index comparison with K-Means algorithm and K-Means++ algorithm is shown in Table 6. In this paper, the FCM clustering algorithm is aimed at the traditional K-Means clustering algorithm, the number of class clusters and the initial cluster center are not well selected, and the repeated calculation of the distance leads to low running efficiency, etc., and chooses the method based on the density and weights to improve the algorithm. Through testing on different behavioral data of students, it is found that the operation results of the algorithm basically match the data distribution of the target dataset. Therefore, this paper applies the algorithm to student behavioral data mining so as to improve the efficiency of data mining to a certain extent.

Table 6. Clustering algorithm evaluation comparison

Algorithm	eval indexing (%)
K-Means	63.11
K-means + +	67.89
This model	77.25

3.2. Student Behavior Correlation Analysis

Correlation coefficients allow for correlation analysis between student behavioral characteristics. The Pearson correlation coefficient was used to analyze the continuity characteristics on the Spark big data platform. The whole process of correlation analysis between student behavioral characteristics is shown in Figure 6.

The correlation coefficient heat map is obtained through the experiment of correlation analysis between student behavioral characteristics. From the figure, it can be seen that students' grades are significantly positively correlated with the number of times students eat breakfast and the number of

books borrowed, i.e., if a student often eats breakfast and borrows books, it is more likely that the student has better academic performance. The number of times a student eats breakfast is correlated with the number of times he/she consumes lunch and the number of times he/she consumes dinner, i.e., if a student eats breakfast, there is a greater likelihood that the student will also eat lunch and dinner.



Fig. 6. Correlation coefficient heat diagram

3.3. Association rule data preprocessing

In order to make the student behavior data satisfy the input format of association rule algorithm, continuous data needs to be transformed into Boolean data.

One semester's data is selected and the continuous student achievement data, book borrowing data and student consumption data are transformed by discretization and clustering techniques. The consumption ability of students is divided into three levels, which are represented by A, B and C, respectively representing low, medium and high consumption levels; the number of breakfasts eaten by students is divided into three levels, which are also divided by FCM algorithm and represented by D, E and F, respectively representing low, medium and high; the number of meals eaten in the cafeteria by students is divided into three levels, which are also divided by FCM algorithm and represented by I, J and K, respectively representing low, medium and high; and the number of meals eaten in the cafeteria by students is divided by I, J and K, respectively representing low, medium and high. The number of students' cafeteria meals is divided into three levels, also using the FCM algorithm to divide them, and represented by I, J, K, respectively, representing low, medium and high; students' performance data are divided into four levels according to the school teaching standards: a for 85-100 points, b for 75-85 points, c for 60-75 points, and d for less than 60 points, representing excellent, medium, passing and failing levels, respectively; students' book borrowing data are discretized by the iso-frequency discretization method, and the book borrowing of less than 3 books is represented by L, and the book borrowing of less than 3 books is represented by L. The student book borrowing data is discretized by equal frequency discretization, with book borrowing less than 3 books expressed as L, book borrowing 3-5 books expressed as M, and book borrowing more than 5 books expressed as H, i.e., low, medium, and high borrowing amount. The results of the preprocessing of some student behavioral data are shown in Table 7.

Table 7. Student behavioral data preprocessing

School number	Consumption capacity	Breakfast times	Dining room times	Average performance	Book borrowing
1	A	F	I	b	H
2	C	D	K	a	L
3	A	F	K	c	M
4	B	E	J	b	H
5	B	E	J	c	M
6	A	F	I	a	M
7	C	D	K	c	L

3.4. Student Behavior Association Rule Mining Analysis

After discrete and clustering transformation of the campus data, the distributed FPGrowth algorithm is used for association rule mining analysis on Spark big data mining platform. The experimental results are shown in Tables 8, 9, and 10. Tables 8 to 10 show the analysis results with minimum confidence level of 0.7, 0.8 and 0.9, respectively.

The distributed FP-Growth association rule algorithm is utilized for mining on Spark big data mining platform. The dataset contains data on student consumption, grades and book borrowing, so the minimum support is set to 0.3. Firstly, the minimum confidence is set to 0.7, and the results are obtained as shown in Table 8, from which it can be seen that when the student's consumption ability is low, the number of times of breakfast consumption and the number of times of cafeteria meals is also low, and the probability of which is 97.18% of the probability; when the students' number of times of meals in the cafeteria is high, the student's consumption ability is medium level, and its possibility is 96.48% probability, it can be found that there is a strong correlation rule between the level of students' consumption ability and the number of breakfast consumption and the number of meals in the cafeteria. When the minimum confidence level is set to 0.8 and 0.9, the probability that the students' grades are also moderate when they have a high number of breakfasts and a moderate number of books borrowed is 96.97% probability. The probability that the students' grades are passing is 91.61% probability when the students have few books borrowed, and it can be found that the students' grades are strongly correlated with the number of breakfasts and the number of books borrowed.

Table 8. Minimum confidence 0.7

Association rule	Confidence
$A \Rightarrow D, I$	0.9718
$A \Rightarrow D$	0.9915
$A \Rightarrow I$	0.9934
$F \Rightarrow C$	0.9488
$K \Rightarrow B$	0.9648
$F, K \Rightarrow C$	0.9911
$c \Rightarrow L$	0.9161
$b \Rightarrow L$	0.8591
$D, L \Rightarrow c$	0.7869
$F, M \Rightarrow b$	0.9697
$F, K \Rightarrow b$	0.896

Table 9. Minimum confidence 0.8

Association rule	Confidence
$A \Rightarrow D, I$	0.9718
$A \Rightarrow D$	0.9915
$A \Rightarrow I$	0.9934
$F \Rightarrow C$	0.9488
$K \Rightarrow B$	0.9648
$F, K \Rightarrow C$	0.9911
$c \Rightarrow L$	0.9161
$F, M \Rightarrow b$	0.9697
$F, K \Rightarrow b$	0.896

Table 10. Minimum confidence 0.9

Association rule	Confidence
$A \Rightarrow D, I$	0.9718
$A \Rightarrow D$	0.9915
$A \Rightarrow I$	0.9934
$F \Rightarrow C$	0.9488
$K \Rightarrow B$	0.9648
$F, K \Rightarrow C$	0.9911
$c \Rightarrow L$	0.9161
$F, M \Rightarrow b$	0.9697

4. Conclusion

This paper designs a multimodal data mining and learning behavior analysis model for Civic Education, and uses improved clustering and association rule algorithms to analyze the multimodal data of students in a college in the sample. It is found that 32.17% of the students in the sample school are above the medium consumption level, 22.08% of the students' average monthly consumption amount fluctuates above and below 378.09 yuan; 10.09% of the students' average monthly consumption amount fluctuates above and below 638.3 yuan. Most of the students are in the below medium level of spending, and the average number of times they spend per month is also low. Only a small number of students have the habit of consuming at school frequently. The student management department of the university needs to appeal to students to develop the habit of consuming at school, which can be done by discounting the price of cafeteria food and improving the quality of cafeteria food to increase the motivation of students to eat at school. In addition, when students eat breakfast a lot and borrow books moderately, the probability that students' grades are also moderate is 96.97%, and there is a strong correlation between students' grades and the number of breakfasts and the number of books borrowed.

References

- [1] Shen, J., Wu, Y., & Liu, W. (2021, May). Research on the Application and Effect of Information Tools in Wisdom Teaching. In *2021 2nd International Conference on Computers, Information Processing and Advanced Education* (pp. 1311-1315).
- [2] Stenberg, K., & Maaranen, K. (2022). Promoting practical wisdom in teacher education: a qualitative descriptive study. *European Journal of Teacher Education*, 45(5), 617-633.

- [3] Yong-jun, Z. H. A. N. G., Ming, Y. A. N. G., Tao, Y. A. N. G., Lan-rong, Z. H. E. N. G., & Yue-e, H. U. A. N. G. (2020). Research process of wisdom classroom and analysis on the hot-spot of knowledge under studies. *Journal of Tropical Diseases and Parasitology*, 18(2), 121.
- [4] Athani, S. S., Kodli, S. A., Banavasi, M. N., & Hiremath, P. S. (2017, May). Student academic performance and social behavior predictor using data mining techniques. In *2017 international conference on computing, communication and automation (ICCCA)* (pp. 170-174). IEEE.
- [5] Huynh, A. C., & Grossmann, I. (2020). A pathway for wisdom-focused education. *Journal of Moral Education*, 49(1), 9-29.
- [6] Wang, R. (2021). Exploration of data mining algorithms of an online learning behaviour log based on cloud computing. *International Journal of Continuing Engineering Education and Life Long Learning*, 31(3), 371-380.
- [7] Wang, R., & Gan, B. (2021). Teaching Activity Model of "PHP Website Development Technology" Course Reform Under the Background of Wisdom Education. In *Application of Intelligent Systems in Multi-modal Information Analytics: Proceedings of the 2020 International Conference on Multi-model Information Analytics (MMIA2020)*, Volume 1 (pp. 252-257). Springer International Publishing.
- [8] AlQaheri, H., & Panda, M. (2022). An education process mining framework: Unveiling meaningful information for understanding students' learning behavior and improving teaching quality. *Information*, 13(1), 29.
- [9] Polizzi, G., & Harrison, T. (2022). Wisdom in the digital age: a conceptual and practical framework for understanding and cultivating cyber-wisdom. *Ethics and Information Technology*, 24(1), 16.
- [10] Arpasat, P., Premchaiswadi, N., Porouhan, P., & Premchaiswadi, W. (2021). Applying process mining to analyze the behavior of learners in online courses. *International Journal of Information and Education Technology*, 11(10), 436-443.
- [11] Chen, L., & Jian, K. (2023, December). Research on Evaluation of Wisdom Education Based on Multi-level Comprehensive Evaluation. In *International Conference on Educational Technology and Administration* (pp. 323-335). Cham: Springer Nature Switzerland.
- [12] Li, S. (2024, September). Research and Development of a Student Behavior Analysis System Based on Big Data Mining. In *2024 3rd International Conference on Artificial Intelligence and Computer Information Technology (AICIT)* (pp. 1-7). IEEE.
- [13] Momay, I. S. I., & Tukang, B. (2023). THE TEACHER'S ROLE IN INTERNALIZING LOCAL WISDOM VALUES AT SMA PGRI KUPANG. *SocioEdu: Sociological Education*, 4(1), 21-26.
- [14] Manzanares, M. C. S., Hernanz, R. J. P., Yáñez, M. J. Z., López, G. A., Sánchez, R. M., Rodríguez, A. C., ... & Arribas, S. R. (2021). Eye-tracking technology and data-mining techniques used for a behavioral analysis of adults engaged in learning processes. *JoVE (Journal of Visualized Experiments)*, (172), e62103.
- [15] Ramdas, G., & Umar, I. N. (2024). MEASURING PRACTICE OF DIGITAL WISDOM IN THE CLASSROOM. *Journal of Educators Online*, 21(3).
- [16] Zhang, J. (2024). Analysis and Modelling Of English Learning Behaviour Based on Data Mining Technology. *Journal of Electrical Systems*, 20(6s), 316-322.
- [17] Albantani, A. M., & Madkur, A. (2018). Think globally, act locally: the strategy of incorporating local wisdom in foreign language teaching in indonesia. *International Journal of Applied Linguistics and English Literature*, 7(2), 1-8.
- [18] Xia, X. (2023). Learning behavior mining and decision recommendation based on association rules in interactive learning environment. *Interactive Learning Environments*, 31(2), 593-608.

- [19] Min, L., Xiaohong, L., & Yan, M. (2022, May). Research on the Blended Teaching Design of Wisdom Classroom under the ADDIE Model. *In 2022 4th International Conference on Computer Science and Technologies in Education (CSTE)* (pp. 87-91). IEEE.
- [20] Gu, Q., Yang, H., & Jiang, N. (2021). Exploration on the Construction of Wisdom Teaching of Universities in the "Internet+" Era. *In Frontier Computing: Proceedings of FC 2020* (pp. 2075-2080). Springer Singapore.
- [21] Yan, S., & Yang, Y. (2021). Education informatization 2.0 in China: Motivation, framework, and vision. *ECNU Review of Education*, 4(2), 410-428.
- [22] Wu, B., & Xiao, J. (2018, October). Mining online learner profile through learning behavior analysis. *In Proceedings of the 10th International Conference on Education Technology and Computers* (pp. 112-118).
- [23] Han, L. (2024, May). Prediction and analysis of students' behavior based on data mining in educational administration. *In International Conference on Artificial Intelligence for Society* (pp. 229-238). Cham: Springer Nature Switzerland.
- [24] Zheng, Z. (2024). Examining The Dual Impact of Comprehensive Psychological Intervention and Ideological and Political Education on College Students' Health Behaviour. *American Journal of Health Behavior*, 48(4), 206-218.
- [25] Liu, J., Wang, C., & Xiao, X. (2021). Internet of things (IoT) technology for the development of intelligent decision support education platform. *Scientific Programming*, 2021(1), 6482088.
- [26] Jie, W., Hai-yan, L., Biao, C., & Yuan, Z. (2017, June). Application of educational data mining on analysis of students' online learning behavior. *In 2017 2nd International Conference on Image, Vision and Computing (ICIVC)* (pp. 1011-1015). IEEE.
- [27] El Aoufi, H., El Hajji, M., Es-Saady, Y., & Douzi, H. (2021). Predicting learner's performance through video sequences viewing behavior analysis using educational data-mining. *Education and Information Technologies*, 26(5), 5799-5814.
- [28] Liu, S., Hu, Z., Peng, X., Liu, Z., Cheng, H. N., & Sun, J. (2017). Mining learning behavioral patterns of students by sequence analysis in cloud classroom. *International Journal of Distance Education Technologies (IJDET)*, 15(1), 15-27.
- [29] Gushchina, O. M., & Ochepovsky, A. V. (2020, May). Data mining of students' behavior in E-learning system. *In Journal of Physics: Conference Series* (Vol. 1553, No. 1, p. 012027). IOP Publishing.
- [30] Chen, L., Wang, L., & Zhou, Y. (2022). Research on data mining combination model analysis and performance prediction based on students' behavior characteristics. *Mathematical Problems in Engineering*, 2022(1), 7403037.
- [31] Khan, A., & Ghosh, S. K. (2021). Student performance analysis and prediction in classroom learning: A review of educational data mining studies. *Education and information technologies*, 26(1), 205-240.
- [32] Zhou, W., & Yang, T. (2021, May). Application analysis of data mining technology in ideological and political education management. *In Journal of Physics: Conference Series* (Vol. 1915, No. 4, p. 042040). IOP Publishing.
- [33] Yang, R. (2021, September). Construction and analysis of ideological and political education reform path based on data mining. *In 2021 4th International Conference on Information Systems and Computer Aided Education* (pp. 214-218).
- [34] Anita Panwar & Satyasai Jagannath Nanda. (2024). Distributed Clustering in Wireless Sensor Network with Kernel Based Weighted Fuzzy C-Means Algorithm. *SN Computer Science*(8),1114-1114.
- [35] Balakrishnan Narayanaswamy,Rychtář Jan,Taylor Dewey & Walter Stephen D..(2024).Accurate approximation of the expected value, standard deviation, and probability density function of extreme

order statistics from Gaussian samples. Communications in Statistics - Simulation and Computation(2),869-878.

- [36] Wu Yilan & Zhang Jing.(2022).Retraction Note: Building the electronic evidence analysis model based on association rule mining and FP-growth algorithm. Soft Computing(1),621-621.