

UI Design Optimization Method Based on User Behavior Analysis - Based on Heat Map Analysis, K-means Clustering, and Random Forest Regression Analysis Algorithms

Yi Yu^{1,✉}, Li Ma², Xiao Chen¹, Yichao Zhong³

¹ HangZhou Animation & Game College, Hangzhou Vocational & Technical College, Hangzhou, Zhejiang, 310000, China

² College of Art, Krirk University, Bangkok, 10220, Thailand

³ New Media Content Center, Hangzhou Bicheng Digital Technology Co., LTD., Hangzhou, Zhejiang, 310000, China

ABSTRACT

In today's digital era, user interface (UI) design is crucial for enhancing user experience and strengthening user engagement. The study uses heatmap analysis, K-means clustering algorithm and random forest regression algorithm to comprehensively analyze the characteristics of user behavior in UI pages. The predicted results of user behavior in UI pages are visualized and analyzed through heatmaps. Cluster classes are divided according to user behavioral characteristics to generate user profiles with the same behavior. Combine Random Forest and Logistic regression algorithm to get the key indexes of UI optimization design and predict their impact on user behavior experience. The research results show that the MAE and SMAPE values of Random Forest regression algorithm on user behavior prediction are 133.55 and 8.18%, respectively, with an R^2 of 0.96, and the accuracy rate of behavior prediction is more than 80%, which shows a good performance of user behavior prediction. The clustering algorithm divides the user behavioral characteristics into 6 clusters based on their behavioral characteristics, including cluster class 1 (browsing and exploring class), which accounts for 11.5% of the number of investigators. The weight of the top 8 of the importance of UI optimization design obtained by the random forest regression analysis algorithm is 70.26%. And the user behavior experience can be improved by 5.377~9.925 times when each element is improved by one unit.

Keywords: Heat map analysis, K-means clustering, Random forest regression, User behavior, UI design

1. Introduction

Heat map analysis is an effective method for visualizing and analyzing key areas in a data set. Through heat map analysis, we can identify regions in a dataset that have a high degree of aggregation and discover relationships and trends between these regions [1-3]. k-means clustering is a classical and widely used clustering algorithm that iteratively calculates the distance between data points and each

✉ Corresponding author.

E-mail address: 2014010010@hzvtc.edu.cn (Y. Yu).

Received 25 March 2024; Revised 05 June 2024; Accepted 15 December 2024; Published Online 15 April 2025.

DOI: [10.61091/jcmcc127a-348](https://doi.org/10.61091/jcmcc127a-348)

© 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/>).

cluster center and classifies the data points into the cluster centers with the closest distance [4-5]. Random forest regression is an integrated learning algorithm based on decision trees, which can effectively cope with the regression problem by constructing multiple decision trees and averaging or voting the prediction results of each decision tree to get the final regression prediction results [6-9]. The above three methods play an important role in the optimization of UI design based on user behavior analysis.

With the rapid development of mobile Internet, users' requirements for websites and applications are getting higher and higher. As a front-end developer, UI design and user experience optimization are very important parts [10-12]. UI design, i.e. user interface design, is the interface design of websites and applications where users interact with the system. A good UI design can not only improve the user experience, but also increase the user's satisfaction with the product [13-16]. Based on user behavior analysis of the use and psychological needs of users, we can better achieve the satisfaction of user needs, thus optimizing the UI, improving user satisfaction, increasing the user retention rate, and ultimately promoting business growth [17-20].

Literature [21] proposed a method to analyze user cognitive and behavioral information in a distributed environment and provide a customized UI, which provides a customized UI by measuring the cognitive abilities and behaviors of users. Literature [22] proposed a collaborative filtering-based method for recommending objects in museums, which effectively simplifies the UI design of museums, improves the recommendation effect, and alleviates the scalability problem. Literature [23] designed a UI framework with personalized and adaptive UI for smart TV users, which is able to change the UI layout and structure according to user needs and contextual details, revealing that it plays an important role in user experience and user satisfaction. Literature [24] examined the acceptance of mHealth-based UI design by older users. It showed that due to its cultural significance, older users are able to accept mHealth-based UI design, and that improving mobile UI design for older users requires considering the cultural rules and behavioral applications of older users. Literature [25] analyzed the dependency relationship of trust on user interface in e-commerce, and by evaluating PGs in marketing campaigns and designing UI for multiple PGs, it revealed that the new UI design is conducive to increasing user stickiness. Literature [26] proposed an evaluation method based on user context and user experience (UX) with a user-centered and statistical evaluation approach, illustrating that the method outperforms existing methods in adapting user interfaces using user context and experience. Literature [27] based on the UI design patterns of marketplace apps, using the Mix method, data collection of users of Shopee, Tokopedia and other apps, and by analyzing the data, a set of design patterns are generated to provide a reference for the research related to the construction of marketplace mobile apps. Literature [28] investigates the role of personalized recommendation systems in improving user experience and emphasizes the importance of behavioral data-driven UI design for business success, revealing that optimized UI design can be achieved through advanced technologies and to user needs to promote effective user interaction. Literature [29] examines the impact of AI-driven UX design on the fintech sector, showing a strong correlation between AI-enhanced features and improved user engagement, providing a reference for AI-driven UX design in fintech. Literature [30] outlines the impact of design and aesthetics on user preferences, i.e., the preference of personality traits for specific interface design features, and discusses the implications of personality-based user interface design, reviewing trends and future directions in the field.

In this study, user behavior patterns in UI pages were explored in depth using a variety of approaches. The study collected data on user interaction behavior, browsing behavior, and decision-making behavior in UI pages. A random forest regression model was used to learn and predict user behavior in UI pages, and the accuracy of user behavior prediction was visualized through heatmaps. Then, K-means clustering algorithm is applied to group the user behavior data and identify user groups with similar behavioral characteristics. Finally, Logistic regression algorithm is applied to predict the relationship between user behavior experience and UI design factors.

2. Research on UI optimization design based on user behavior

2.1. Optimization goals in UI design

The optimization goals in UI design [31] mainly include three aspects of visual design, interaction design, and emotional design, which are detailed below:

2.1.1. Visual design. The design of visual information helps users to acquire information in an orderly manner and process it logically. Although visual design has lost its original position in the whole interface design process, it is not yet possible to reach the era of interface without visual design. Three elements of information acquisition are involved in the process of acquiring information:

- 1) The quality of the acquired information: it refers to the accuracy of the visual need to convey information, which involves the text, graphics, colors and other elements and their combination after various connections.
- 2) Quantity of acquired information: It refers to the total number of letters visually transmitted.
- 3) The speed of information processing: the process of information processing is the process of cognition. The speed of information processing depends on the speed of cognition, which is not closely related to the speed of information acquisition. If visual information is grouped and logical, it can achieve the effect of recognizing instead of memorizing and then reacting.

2.1.2. Interaction design. In this paper, we believe that the optimization goal of UI interaction design will be to elevate interaction design to the goal of design planning: interaction design will be more targeted according to the user's needs and goals for the realization of the program. It will have a strong logic, become a marketing tool, and also serve as a platform for communication between designers and users.

- 1) Targeted. Refers to the design before a thorough research is necessary to analyze the user's needs, goals, and make a detailed, detailed plan.
- 2) Logic. It is a product of the increased status of interaction design. The improvement of logic is conducive to leading the whole software design process, becoming the blueprint of the whole process, and leading other parts of the whole production process with its own description. In addition, from the user's point of view, a logical interaction process can guide the user to ease of use and effectively help them accomplish their current tasks. Thus logic is increasing in all aspects of interaction design.
- 3) Marketing Tools. The process of interaction is the process of communication, in which it can help marketers to understand the motivation of users to use the product, and according to the data statistics, orientation analysis and other steps can make the marketing process more inclined.
- 4) A platform for designers to communicate with users. Designers must combine the use with research, and let users provide their own experience to become the participants of design and update in the interaction design. And this can make the designer understand the user better, back to the first link, so that the whole design process becomes a virtuous cycle.

2.1.3. Emotional design. Emotional design is an inevitable factor in UI design, however generation to the cognitive process can be found:

Any emotional fluctuation of the user will sway the user's perception of the interface, thus making them react differently. There is a clear intersection between emotional design and visual design and interaction design, and every design may contain emotional factors. If an interface has an emotional reaction of disgust to its messy arrangement at the first sight of the user, then the whole process of using it will carry this feeling of negative emotion, and thus will make a poor or wrong overall perception.

This paper argues that although it is impossible to know the specific personal factors of the user, it is better to pre-design the emotions that need to be generated in the process of using the interface rather than letting the user's emotions dominate the process of using it.

Emotional overlay is the simplest step in emotional design, only need to add some elements that can guide the emotions in the interface design itself to achieve the mouth. Such as with ghosts and gory images can make the user produce fear, nausea, this emotion and the design of the content in line with the user's emotions will be superimposed, so that the interface will remain in the memory after use.

2.2. *Application of user behavior in UI design*

After analyzing the user behavior of the mobile page and coming up with the reasons, the focus and user behavior can be applied to the UI design to make it more in line with the user experience.

1) Too much data on the mobile page will lead to slow loading and ultimately lead to user loss. That is the focus of the designer must regulate the data, too much data will make the single page open slowly, flat and light style popularity is also to solve such problems and sprouted. The difference between flat icons and anthropomorphic ones is the elimination of some superfluous stylized effects. The important thing is that its low data, the loading of the page has a good effect on reducing the burden.

2) for the mobile page click peak period can be seen, are entertainment or rest time, then mobile developers in the push notification should try to choose this period, in order to maximize the information publicity, and version updates and other situations that will lead to software can not browse the situation as far as possible to be arranged with the trough, so that according to the user's habit of browsing time to arrange the various types of functions in order to avoid the shortcomings of the strengths and improve the user experience.

3) The heat of the page for 48 hours, which is the same applied to the mobile APP new information push, around the user's attention to adjust and version iteration push time interval, too frequent push may cause information congestion to the user and the increase in the degree of omission.

4) most users right hand for the habitual hand, which should be through the layout of the software layout to cope with the frequent click on the function area should be set in the user easy to touch the position, and easy to misoperation, such as return, delete, etc. should be set in the user is not easy to touch or need to adjust the posture of the click on the area, and set up left and right sliding into and out of the function is in line with the user's behavior to improve the user experience of the design.

5) The reason for the low forwarding rate of the mobile page, a large part of it is used for the content, under the premise of complete and excellent content, it is extremely important to do a good job of forwarding the recipient's perfect, WeChat, QQ, microblogging and other mainstream social platforms are indispensable.

3. **Key technologies and methodologies**

3.1. *Heat map analysis*

A heat map [32] is a technique in which each individual value contained in a numerical matrix is graphically represented by a color. Heat maps originated to display values in a two-dimensional matrix of values, with larger values represented by dark gray or black squares (pixels) and smaller values by lighter colors. Heatmaps illustrate the distribution of discrete data (events or things) and their interrelationships, often presenting the end result as a picture with significant color differences.

Heatmaps are a popular data visualization method. The data matrices presented are usually obtained by rearranging the rows and columns of the matrix using a specific algorithm. Heatmap data sources are various statistical analysis tools, soft armor, computer scripting languages and so on. The two dimensions of the matrix columns of the heat map can be the same category of things, the data in the matrix reflects the interconnection between similar things. It is also possible that the rows and columns are things and different conditions, and the data in the matrix reflect the changes in the properties of things in different states.

3.2. *K-means clustering*

K-means algorithm [33] is an unsupervised clustering algorithm based on partitioning. The K of K-means algorithm is the number of clusters, which divides all the samples into K classes so that the distance between the classes is as large as possible, and the distance within the classes is as small as possible. The K-means algorithm is generally divided into two steps: first, iterates over each sample, and determines the point closest to the center of all centroids as its center. Then, according to the new center distribution of each sample to re-determine the clustering center point. The above steps are repeated until the centroids of each sample no longer change.

3.2.1. Similarity measures for clustering algorithms. General clustering algorithms finalize the clustering of samples by measuring the similarity between samples and performing the merging of similar samples. Clustering algorithms can use different distance functions or similarity functions to measure the relationship between samples depending on the usage scenario.

1) Euclidean distance. Euclidean distance is one of the more widely used similarity measure functions, and in this paper, Euclidean distance is selected as the similarity measure function:

$$\left(\sum_{l=1}^d |x_{il} - x_{jl}|^2 \right)^{1/2} \quad (1)$$

2) Cosine distance. Cosine distance is one of the most commonly used similarity metric functions in the field of text categorization, which can remain constant for rotational changes in the data:

$$1 - \cos \alpha = \frac{x_i^T x_j}{\|x_i\| \|x_j\|} \quad (2)$$

3) Pearson correlation distance. Pearson correlation distance measures the similarity between sample points based on linear correlation:

$$1 - \frac{\text{Cov}(x_i, x_j)}{\sqrt{D(x_i)} \sqrt{D(x_j)}} \quad (3)$$

where Cov denotes covariance and D denotes variance.

3.2.2. Metrics for clustering algorithms. In this paper, contour coefficient (Si) and Davies-Bouldin (DB) are used as the measures of clustering algorithm:

1) Contour coefficient. The contour coefficient evaluates the clustering results by calculating the average distance between the sample points and other sample points in the same class, as well as the average distance between different clusters. Each class is represented by a contour plot, and the contour combines the intra-class dissimilarity and inter-class dissimilarity of the clusters, which visualizes which sample points are located within the class clusters. Based on the intra-class dissimilarity and inter-class dissimilarity of the sample points the contour coefficient of the sample points can be calculated:

$$S_i = \frac{b_i - a_i}{\max\{a_i, b_i\}} \quad (4)$$

where a_i is the within-class dissimilarity of the sample point, which represents the average distance of the sample point t from other samples of the same class cluster, and $b_i = \min\{b_{i1}, b_{i2}, \dots, b_{ik}\}, j=1, 2, \dots, k, b_i$ is the between-class dissimilarity of the sample point, which represents the minimum of the average distance b_{ij} of the sample point t from all other samples of a certain class cluster C_j .

2) Davies-Bouldin Indicator. The Davies-Bouldin metric indicates the ratio of the intraclass scatter of

a sample point to the distance between the cluster centroids:

$$DB = \frac{1}{k} \sum_{i=1}^k \max_{i \neq j} \left(\frac{\sigma_i + \sigma_j}{d(c_i, c_j)} \right) \quad (5)$$

where k denotes the number of clusters, c_x denotes the centroid of class x , σ_x denotes the average distance between all sample points in class x and centroid c_x , and $d(c_i, c_j)$ denotes the distance between centroid c_i and c_j , where the distance calculation is generally done using the Euclidean distance, which is more limited for streaming and high-dimensional data.

3.3. Random Forest Regression Analysis

The Random Forest model is derived from the CART Decision Tree model, combining the Bagging integrated learning method and random subspace theory. The random forest model is capable of handling both discrete and continuous variables.

3.3.1. CART regression decision tree. CART decision tree [34] is a binary tree structure with nodes divided into leaf nodes and non-leaf nodes. At each non-leaf node a feature attribute is selected as a split variable, and two different branches are constructed based on the division of this split variable. When the samples under a node cannot or need not be further divided, the node is a leaf node, and each leaf node stores a value or a category.

Specifically, for the regression decision tree model, assume that the training sample T has n characteristic attributes $x_1, x_2, \dots, x_n$. If the variable x_i at the non-leaf node is a continuous variable, and its split point is set to be s_i , then the sample space is divided into two parts, $\{X | x_i \leq s_i\}$ and $\{X | x_i > s_i\}$. Split point s_i is determined by traversing all possible values of x_i , and the criterion for determining it is to minimize the sum of the mean squared error (MSE) of each of the two regions, and following this rule, assuming that the final sample space is partitioned into m regions $R_1, R_2, \dots, R_m$, the regression tree has m leaves, each of which corresponds to a predicted value. Let the j th leaf store the predicted value c and c_j be the mean of the true values of the samples belonging to the space R_j .

3.3.2. Steps in the Random Forest Regression Algorithm. The working principle of the random forest regression algorithm is shown in Figure 1, and the specific steps [35] are as follows:

1) The bootstrap method is used to randomly draw K training samples from the original samples with putback to generate K decision trees, and the size of each training sample is the same as the size of the original sample. Assuming that there are N samples in the original set of samples, each round of extraction is randomly drawn N times with putative randomization. The probability that each sample will not be drawn in a round of extraction is $(1 - 1/N)^N$. When N is large enough, this probability converges to $1/e$, which means that for each round of extraction, about 36.8% of the original samples will not be drawn.

2) Each decision tree grows independently under its respective training set, with rows branching at non-leaf nodes according to the aforementioned principle of minimizing the sum of mean square errors. According to the method provided by the random subspace theory, when generating each non-leaf node, m variables are randomly selected from all the feature variables to form a subset of feature variables, and then the one feature variable with the most regression effect is selected from this subset as the split node. Assuming that the sample has a total of M feature attributes, in the random forest regression algorithm, usually take $m = M / 3$.

3) Each decision tree grows independently to get different branching structures. When making decisions, the input data will pass through all K decision trees and get K regression results, and the final prediction of the model is the arithmetic mean of the respective regression results of all decision trees. Assuming that input data x passes through the i th decision tree to obtain a regression result

of $h_i(x)$, then the prediction result of the regression to the random forest is:

$$H(x) = \frac{1}{K} \sum_{k=1}^K h_k(x) \quad (6)$$

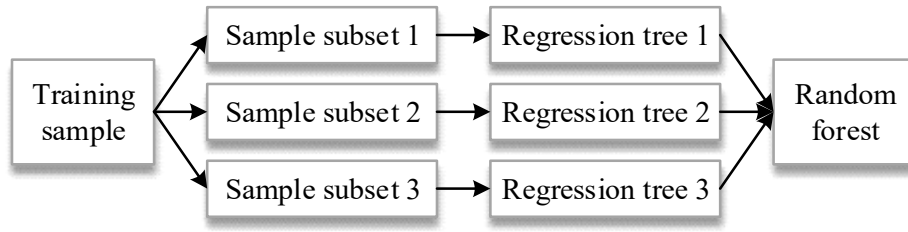


Fig. 1. Schematic diagram of how a random forest works

3.3.3. Indicators of variable importance:

1) Importance scoring based on Gini coefficient. Let the characteristic attribute of sample T in a decision tree at non-leaf node m be x_i and the probability that the sample belongs to category j be p_j , then the Gini coefficient of variable x_i at node m is:

$$Gini(T) = 1 - \sum p_j^2 = 1 - \sum \left(\frac{n_j}{S}\right)^2 \quad (7)$$

Where S is the size of the sample T , and n_j is the number of samples belonging to class j in the sample T . The larger the Gini coefficient is, the less “pure” the node is. The larger the Gini coefficient, the lower the “purity” of the node, that is to say, the more dispersed the sample data at the node.

Assuming that sample T splits into T_1, T_2 two samples through two branches at node m , and the sample size is s_1, s_2 , then the importance statistic of variable x_i at node m is:

$$VIM_{i,m}^{Gini} = Gini(T) - \frac{s_1}{S} Gini(T_1) - \frac{s_2}{S} Gini(T_2) \quad (8)$$

Let variable x_i occurs M times in this decision tree, then the importance statistic of variable x_i in this tree is:

$$VIM_i^{Gini} = \sum_{m=1}^M VIM_{i,m}^{Gini} \quad (9)$$

Assuming that there are K decision trees in this random forest, the importance score of variable x_i in this random forest is:

$$VIM_i^{Gini} = \frac{1}{K} \sum_{k=1}^K VIM_{i,k}^{Gini} \quad (10)$$

VIM The higher the score, the higher the importance of the variable.

2) Importance scores based on prediction errors of out-of-bag data. From the bootstrap random sampling method described in the previous section, it can be seen that after the generation of the random forest, for each decision tree in the random forest, there are about 36.8% of the samples in the original sample set have not been sampled as its training set, and these samples are known as the “out-of-bag (OOB) data”.

By regressing the out-of-bag data in decision tree k , the mean square error of this regression result can be obtained MSE_k . Randomly replacing the values of the explanatory variable x_i and regressing it with the same decision tree yields the mean square error of the new regression result MSE_{ik} . Normalizing the difference between these two mean square errors yields the significance statistic of variable x_i in the k th decision tree $VIM_{i,k}^{OOB}$. If $VIM_{i,k}^{OOB}$ is 0, it means that variable x_i does not

occupy any node in the k th decision tree.

Assuming that there are K decision trees in the random forest, the importance score of explanatory variable x_i in this random forest is:

$$VIM_i^{OOB} = \frac{1}{K} \sum_{k=1}^K VIM_{i,k}^{OOB} \quad (11)$$

The higher the score, the greater the impact of the change in the value of the explanatory variable on the regression results, representing a higher level of importance of the variable.

4. UI design optimization analysis under user behavior analysis

Experimental environment: Intel(R)Core(TM) i3-7100U CPU@3.90 GHz. Software environment: Tenchi Notebook. In this paper, we collected a large amount of user behavior data in the UI of online platforms, which contains a large amount of user behavior data such as clicking, swiping, and typing for the following experimental study.

4.1. Heat map analysis of user behavior prediction results

In the prediction of user purchase behavior, the error analysis and evaluation of the model is an important link, and its evaluation index can comprehensively reflect the working status of the system. The following evaluation criteria are usually used to analyze the prediction performance of the model for the sample:

- 1) MAE stands for Mean Absolute Error, which measures the average of the absolute error between the predicted value and the true value, the smaller the value of MAE, the more accurate the prediction of the model.
- 2) R^2 , which represents the fit of the model to the data, i.e., the ability of the model to explain the variability of the observations. The value of R^2 is usually between 0 and 1, and may be negative, the larger the value, the more accurate the prediction of the model.
- 3) SMAPE, i.e. Symmetric Mean Absolute Percentage Error. The range of SMAPE is from 0% to 100%, and a lower SMAPE value indicates a higher prediction accuracy.

The formula for the evaluation index is as follows:

$$MAE = \frac{1}{n} \sum_{i=1}^n |y_i - \hat{y}_i| \quad (12)$$

$$R^2 = 1 - \frac{\sum_i (y_i - \hat{y}_i)^2}{\sum_i (\bar{y}_i - y_i)^2} \quad (13)$$

$$SMAPE = \frac{100\%}{n} \sum_{i=1}^n \frac{|\hat{y}_i - y_i|}{(|\hat{y}_i| + |y_i|)/2} \quad (14)$$

Where, n denotes the number of predictions, y_i denotes the true value, $\hat{y}_i$ denotes the predicted value, and $\bar{y}_i$ denotes the mean value.

The SVM, MLP, LSTM, CNN, and BP models were also compared with the random forest regression analysis algorithm in this paper to verify the effectiveness of the models. The different models were evaluated using MAE, R^2 , SMAPE, and the final quantitative computational error results are shown in Figure 2. As can be seen from the figure, the random forest model used in this paper has significantly improved the prediction accuracy of user behavior, the MAE value and SMAPE value show the lowest 133.55 and 8.18% respectively, while R^2 is close to 1, 0.96. It shows that the random forest algorithm can effectively achieve the prediction study of user behavior in the UI page.

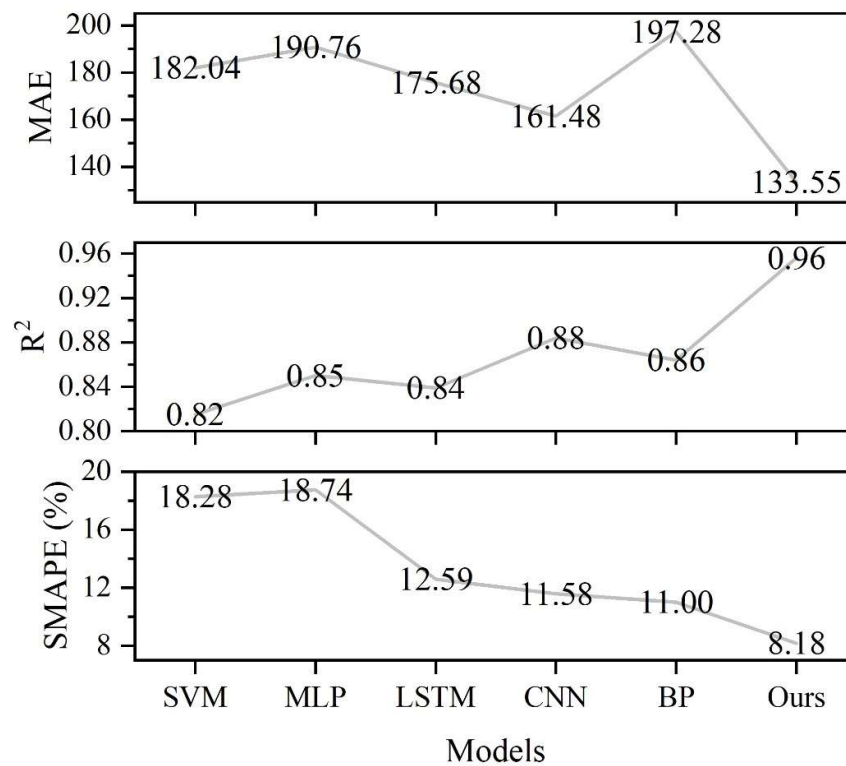


Fig. 2. Final quantitative calculation error

In order to further analyze the results of the Random Forest algorithm for predicting user behaviors on UI pages, heatmaps are used to observe the performance of the prediction model proposed in this paper on each behavioral category. User behaviors include interaction behaviors (clicking, swiping, typing, dragging), browsing behaviors (scanning, reading), and decision-making behaviors (selecting, confirming).

The results of user behavior prediction for UI pages based on the random forest algorithm are shown in Figure 3. The data in the figure shows the frequency of accurate prediction of each user behavior by the random forest algorithm, and it can be seen that the accuracy of user behavior prediction of UI page using the random forest algorithm is high. After calculation, the average prediction accuracy of user interaction behavior, browsing behavior, and decision-making behavior are 86.06%, 87.00%, and 86.54%, which is not a big difference. In addition, in terms of individual user behaviors, the accuracy rate of random forest prediction is 83.98% to 89.10%, with the confirmation behavior in decision-making behavior having the highest accuracy rate. It indicates that the random forest model achieves a better balance between user behavior bias and variance by adjusting the hyperparameters such as the number of decision trees, the depth of the number, and the proportion of feature selection, so that the model achieves a better balance between user behavior bias and variance. The prediction results can meet the reference of UI optimization design based on user behavior.

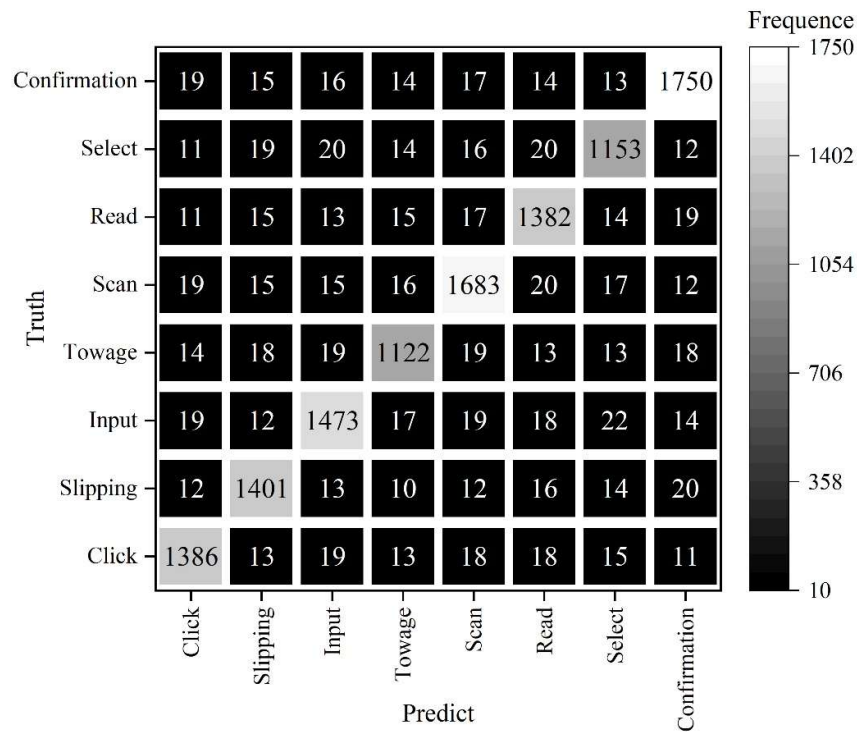


Fig. 3. The user behavior prediction results based on the random forest algorithm

4.2. User Cluster Classification Based on Behavioral Characteristics

In this section, based on the experiments described above, 200 users in an online platform were extracted, and a clustering analysis based on user behavioral characteristics was performed on their eight behavioral characteristics of clicking, sliding, typing, dragging, scanning, reading, selecting and confirming in the UI.

The clustering effect of the K-means clustering algorithm is evaluated using the clustering assessment indexes introduced above (Si, DB index), and the results of different numbers of clusters Si and DB indexes are shown in Figure 4. From the figure, it can be seen that when the number of clusters of clustering effect is 6, the Si value is the largest of about 0.755, and the DB value shows the smallest, and the clustering effect is the best. It shows that the K-means clustering algorithm is able to carry out preliminary clustering of UI user value subgroups to further analyze the user behavior in the UI of online platforms.

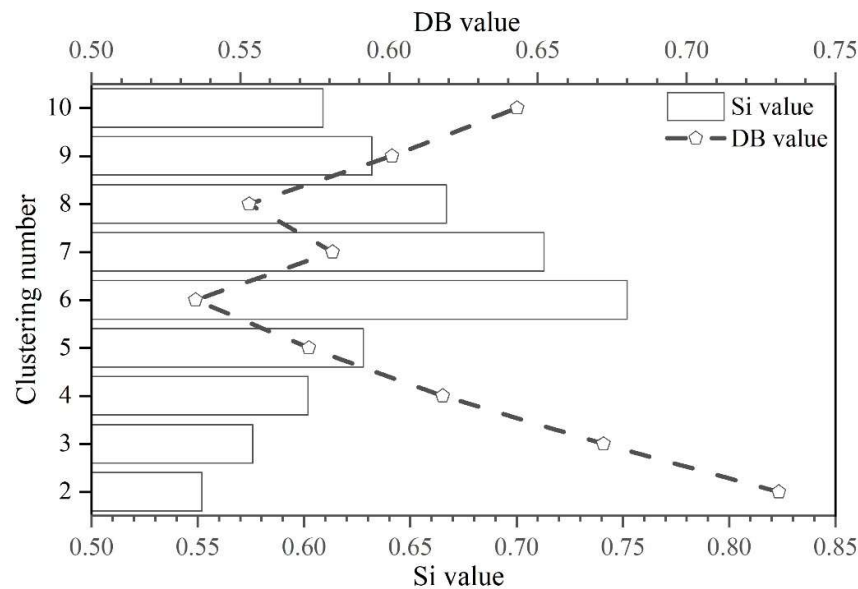


Fig. 4. The results of the SI and DB indices of different clustering Numbers

Through the determination of the optimal number of clusters above, this section obtains the clustering division results based on user behavioral characteristics as shown in Fig. 5. From the figure, it can be seen that the user behavioral characteristics are divided into six clusters, each of which is uniformly distributed around the respective cluster center, and the data points between the clusters are relatively dispersed. The six clusters are: browsing and exploring (Cluster 1), precise operation (Cluster 2), interaction and communication (Cluster 3), creation and generation (Cluster 4), task execution (Cluster 5), and input and search (Cluster 6), respectively. (Cluster 6). Cluster 1 is taken as an example for specific analysis:

Cluster 1 (Browse and Explore), accounting for 11.5% of the survey respondents, users in this cluster frequently clicked on different areas of the UI page to view different content. And they continuously slide the UI page, less often use drag-and-drop behavior, and occasionally drag the scroll bar to quickly locate the page position. In addition, this type of user inputs less content and mainly searches for keywords. By quickly scanning the surface to browse the UI page layout, titles, images and other elements to find the information points of interest. Their reading content is shallow, mainly short descriptions such as titles and summaries. Selection is made according to the guidance of the first image or UI page layout, and the confirmation behavior is mainly a simple confirmation of the search results and browsing content.

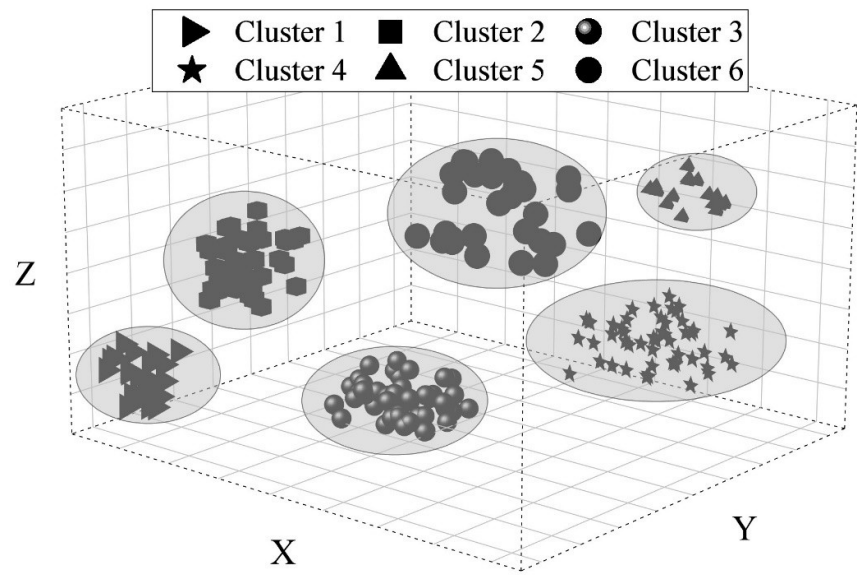


Fig. 5. Clustering of clustering based on user behavior characteristics

4.3. Evaluation of the importance of UI optimization design elements

In order to deeply explore the importance of each element in UI optimization design, this section, through literature review and consulting with UI design related experts, designs the element indicators in UI design as shown in Table 1. Specifically, it includes: layout (layout partition, layout to its, layout white space), color (color matching, color symbols, color guidance), font (font selection, font size and font weight, font spacing), icon (icon design, size and resolution, icon color), interaction (operation feedback, operation flow, fault-tolerant mechanism) and emotional design.

Table 1. Elements and number of UI design

Macroelements	Specific elements	Number
Layout design	Layout partition	N1
	Layout pair	N2
	Layout white	N3
Color design	Color collocation	N4
	Color symbol	N5
	Color guide	N6
Font design	Font selection	N7
	Word size and word weight	N8
	Word spacing	N9
Icon design	Icon design	N10
	Dimension and resolution	N11
	Icon color	N12
Interaction design	Operational feedback	N13
	Operating procedure	N14
	Fault tolerant mechanism	N15
Emotional design	Emotional design	N16

The study obtained the importance weights of the elements through the Random Forest algorithm to show the importance of the elemental indicators to the UI optimization design, and Figure 6 demonstrates the weight values of the elements in the UI optimization design. Among them, N16 (emotional design) contributes the most to UI optimization design, with element weights of 11.5%.

Through emotional design, users can feel pleasure and comfort when using the product, thus increasing user satisfaction. It is followed by N4 (color matching) with a weight of 10.7%, and N13 (operation feedback) in third place with a weight of 10.0%. The fourth to eighth most important design elements and their weights are N10 (icon design, 9.3%), N5 (color symbols, 8.2%), N11 (size and resolution, 7.5%), N15 (fault-tolerance mechanism, 6.7%), and N3 (layout white space, 6.35%). The total weight of the above 8 elements is 70.26%. The weights of the last 8 items, in descending order, are N7 (font selection), N9 (font spacing), N8 (font size and font weight), N1 (layout partitioning), N6 (color guidance), N12 (icon color), N14 (operation flow), and N2 (layout to its), with the weights of 1.59%, 2.2%, 3.1%, 3.6%, 4.03%, 4.38%, 5.23% and 5.61%. By analyzing the importance of elements, key indicators with important impact on UI design optimization are identified to provide a reasonable basis for the development of UI design optimization methods.

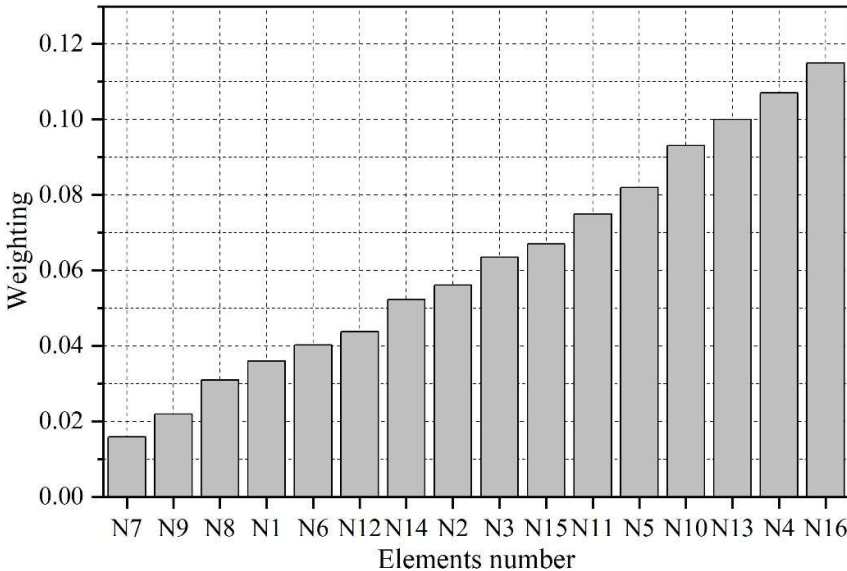


Fig. 6. The weight of each element in the UI optimization design

4.4. Logistic Regression Analysis of User Behavior Experience

In this section, based on the Random Forest algorithm, the UI optimization design influencing factors with the highest importance ratings are included in the multifactor Logistic regression analysis model to assess and analyze the role and effect values of UI design optimization influencing factors on user behavioral experience. The selected factors include the top 8 elements of importance, namely N16, N4, N13, N10, N5, N11, N15 and N3. The UI user behavior experience is taken as the dependent variable.

Table 2 shows the results of logistic regression analysis of the relevant elements of UI optimization and design. As can be seen from the table, the overall significance of the model has a P-value of 0.000 less than 0.05, indicating that the model is statistically significant in general, and the P-value of the Hosmer-Lemeshow test is 15.496 greater than 0.05, which means that the model is well fitted. The regression coefficients of each factor in UI optimization and design range from 0.113 to 0.282, and the P-values are all 0.000, indicating that each factor has a positive influence on user behavioral experience and passes the significance test at 0.001 level. In addition, according to the OR value analysis in the table, the OR value of the influencing factors is between 5.377~9.925, which is greater than 1, indicating that for each unit of elevation of each influencing factor, the chances of the occurrence of the event of the dependent variable increase by 5.377~9.925 times.

Table 2. Analysis of logistic regression analysis of UI optimization design factors

Variable	Regression coefficient	SE	P	OR	95%CI
N16	0.282	0.030	0.000***	9.925	3.36~6.95
N4	0.187	0.033	0.000***	9.242	3.333~5.702
N13	0.348	0.053	0.000***	8.854	3.13~6.743
N10	0.113	0.021	0.000***	8.544	2.314~4.839
N5	0.148	0.034	0.000***	8.129	3.163~6.987
N11	0.264	0.011	0.000***	5.927	3.412~6.791
N15	0.245	0.037	0.000***	5.859	2.758~4.858
N3	0.241	0.025	0.000***	5.377	2.351~4.428
Overall significance test		P=0.000			
Hosmer-Lemeshow		P=15.496			

Note: ***, indicates that the p-value passes the test of significance at the 0.001 level.

5. Conclusion

The article analyzes the optimization goals in UI design and the application of user behavior in UI design, and launches a research on UI design optimization based on user behavior analysis. The study uses heatmap analysis to visualize the results of Random Forest regression analysis algorithm for predicting user behavior on UI pages. K-means clustering algorithm identifies different user behavior patterns and groups. Random forest and logistic regression analysis algorithms are further used to predict user behavior experience.

- 1) Random Forest regression analysis algorithm has better user behavior prediction performance, its MAE value and SMAPE value are significantly lower than the comparison method, and the accuracy of user behavior prediction on UI page is between 83.98% and 89.10%.
- 2) In this paper, the user behavior is divided into 6 cluster classes based on their characteristics, and the Si and DB values of the clustering effect evaluation indexes are about 0.755 and 0.535, respectively.
- 3) The study obtains the most critical index in UI optimization design according to Random Forest regression algorithm as emotional design with a weight of 11.5%.
- 4) It is obtained by Logistic regression algorithm that each UI optimization design element has better improvement for user behavior experience.

This study provides data-driven decision support for UI optimization design, optimizes product design, improves user experience and drives business decisions.

References

- [1] Pati, S. K., Banerjee, A., & Manna, S. (2023). Gene selection of microarray data using heatmap analysis and graph neural network. *Applied Soft Computing*, 135, 110034.
- [2] Hong, I., & Jung, J. K. (2017). What is so “hot” in heatmap? Qualitative code cluster analysis with foursquare

- venue. *Cartographica: The International Journal for Geographic Information and Geovisualization*, 52(4), 332-348.
- [3] Rutgers, D. R., Van der Gijp, A., Vincken, K. L., Mol, C. P., Van der Schaaf, M. F., & Ten Cate, T. J. (2021). Heat map analysis in radiological image interpretation: an exploration of its usefulness for feedback about image interpretation skills in learners. *Academic Radiology*, 28(3), 414-423.
 - [4] Sinaga, K. P., & Yang, M. S. (2020). Unsupervised K-means clustering algorithm. *IEEE access*, 8, 80716-80727.
 - [5] Yuan, C., & Yang, H. (2019). Research on K-value selection method of K-means clustering algorithm. *J*, 2(2), 226-235.
 - [6] Coulston, J. W., Blinn, C. E., Thomas, V. A., & Wynne, R. H. (2016). Approximating prediction uncertainty for random forest regression models. *Photogrammetric Engineering & Remote Sensing*, 82(3), 189-197.
 - [7] Borup, D., Christensen, B. J., Mühlbach, N. S., & Nielsen, M. S. (2023). Targeting predictors in random forest regression. *International Journal of Forecasting*, 39(2), 841-868.
 - [8] Wang, F., Wang, Y., Zhang, K., Hu, M., Weng, Q., & Zhang, H. (2021). Spatial heterogeneity modeling of water quality based on random forest regression and model interpretation. *Environmental Research*, 202, 111660.
 - [9] Rigatti, S. J. (2017). Random forest. *Journal of Insurance Medicine*, 47(1), 31-39.
 - [10] Punchoojit, L., & Hongwarittorn, N. (2017). Usability studies on mobile user interface design patterns: a systematic literature review. *Advances in Human-Computer Interaction*, 2017(1), 6787504.
 - [11] Tsai, T. H., Chang, H. T., Chen, Y. J., & Chang, Y. S. (2017). Determinants of user acceptance of a specific social platform for older adults: An empirical examination of user interface characteristics and behavioral intention. *PloS one*, 12(8), e0180102.
 - [12] Arazy, O., Nov, O., & Kumar, N. (2015). Personalization: UI personalization, theoretical grounding in HCI and design research. *AIS Transactions on Human-Computer Interaction*, 7(2), 43-69.
 - [13] Orlovskaya, J., Wickman, C., & Söderberg, R. (2018). Big Data Usage Can Be a Solution for User Behavior Evaluation: An Automotive Industry Example. *Procedia CIRP*, 72, 117-122.
 - [14] Chen, J., Chen, C., Xing, Z., Xia, X., Zhu, L., Grundy, J., & Wang, J. (2020). Wireframe-based UI design search through image autoencoder. *ACM Transactions on Software Engineering and Methodology (TOSEM)*, 29(3), 1-31.
 - [15] Miraz, M. H., Excell, P. S., & Ali, M. (2016). User interface (UI) design issues for multilingual users: a case study. *Universal Access in the Information Society*, 15, 431-444.
 - [16] Johnson, J. (2020). *Designing with the mind in mind: simple guide to understanding user interface design guidelines*. Morgan Kaufmann.
 - [17] Ruiz, J., Serral, E., & Snoeck, M. (2021). Unifying functional user interface design principles. *International Journal of Human-Computer Interaction*, 37(1), 47-67.
 - [18] Sridevi, S. (2014). User interface design. *International Journal of Computer Science and Information Technology Research*, 2(2), 415-426.
 - [19] Miraz, M. H., Ali, M., & Excell, P. S. (2021). Adaptive user interfaces and universal usability through plasticity of user interface design. *Computer Science Review*, 40, 100363.
 - [20] Nguyen, T. T., Vu, P. M., Pham, H. V., & Nguyen, T. T. (2018, May). Deep learning UI design patterns of mobile apps. In *Proceedings of the 40th International Conference on Software Engineering: New Ideas and Emerging Results* (pp. 65-68).

- [21] Ji, H., Yun, Y., Lee, S., Kim, K., & Lim, H. (2018). An adaptable UI/UX considering user's cognitive and behavior information in distributed environment. *Cluster Computing*, 21, 1045-1058.
- [22] Zhang, L., & Gao, Y. (2022). UI Design and Optimization Method for Museum Display Based on User Behavior Recommendation. *Wireless Communications and Mobile Computing*, 2022(1), 2814216.
- [23] Khan, M., & Khusro, S. (2023). Towards the design of personalized adaptive user interfaces for smart TV viewers. *Journal of King Saud University-Computer and Information Sciences*, 35(9), 101777.
- [24] Alsswey, A., & Al-Samarraie, H. (2020). Elderly users' acceptance of mHealth user interface (UI) design-based culture: the moderator role of age. *Journal on multimodal user interfaces*, 14, 49-59.
- [25] Yashmi, N., Momenzadeh, E., Anvari, S. T., Adibzade, P., Moosavipoor, M., Sarikhani, M., & Feridouni, K. (2020, May). The effect of interface on user trust; user behavior in e-commerce products. In *Proceedings of the Design Society: DESIGN Conference (Vol. 1, pp. 1589-1596)*. Cambridge University Press.
- [26] Hussain, J., Ul Hassan, A., Muhammad Bilal, H. S., Ali, R., Afzal, M., Hussain, S., ... & Lee, S. (2018). Model-based adaptive user interface based on context and user experience evaluation. *Journal on Multimodal User Interfaces*, 12, 1-16.
- [27] Arifin, J., Fasha, E. M., & Ayu, M. A. (2021, August). User Interface Design Patterns for Marketplace Mobile Application in Indonesia Based on User Behavior. In *2021 IEEE 7th International Conference on Computing, Engineering and Design (ICCED)* (pp. 1-6). IEEE.
- [28] Liu, Y., Xu, Y., & Zhou, S. (2024). Enhancing User Experience through Machine Learning-Based Personalized Recommendation Systems: Behavior Data-Driven UI Design. *Authorea Preprints*.
- [29] Xu, Y., Liu, Y., Xu, H., & Tan, H. (2024). AI-driven UX/UI design: Empirical research and applications in FinTech. *International Journal of Innovative Research in Computer Science & Technology*, 12(4), 99-109.
- [30] Alves, T., Natálio, J., Henriques-Calado, J., & Gama, S. (2020). Incorporating personality in user interface design: A review. *Personality and Individual Differences*, 155, 109709.
- [31] Ya Yu, Fangjun Liu & Yuan He. (2024) .Teaching Reform Design of UI Design Course Based on BOPPPS Teaching Model. *International Journal of New Developments in Education*(11).
- [32] Rujia Chen, Akbar Ghobakhlou & Ajit Narayanan. (2024). Interpreting CNN models for musical instrument recognition using multi-spectrogram heatmap analysis: a preliminary study. *Frontiers in Artificial Intelligence* 1499913-1499913.
- [33] Mingchao Qi, JunQiang Zhao & Yan Feng. (2024). An optimized public opinion communication system in social media networks based on K-means cluster analysis. *Heliyon*(24), e40033-e40033.
- [34] Elsayed A. Abdelsamie, Abdel rahman A. Mustafa, Abdelbaset S. El Sorogy, Hanafey F. Maswada, Sattam A. Almadani, Mohamed S. Shokr... & Jose Emilio Meroño de Larriva. (2024). Current and Potential Land Use/Land Cover (LULC) Scenarios in Dry Lands Using a CA-Markov Simulation Model and the Classification and Regression Tree (CART) Method: A Cloud-Based Google Earth Engine (GEE) Approach. *Sustainability*(24), 11130-11130.
- [35] Wang Zhengqing, Yu Yan, Chen Yulin, Liu Zhongnan, He Haiyang, Guo Zhixin... & Pan Junhua. (2025). Geohazard Sensitivity Evaluation in Xinning, Hunan, China, Using Random Forest, Artificial Neural Network, and Logistic Regression Algorithms. *Natural Hazards Review*(2).