

A study on the application of AI algorithms in cross-cultural marketing strategies

Mingyue Liu¹, 

¹ Department of Business and Commerce, Zhengzhou Vocational College of Finance and Taxation, Zhengzhou, Henan, 450048, China

ABSTRACT

With the development of globalization, the cross-cultural market is facing needs such as diversification and personalization of consumer demand. Based on the theory of market segmentation, the study proposes an ant colony algorithm to improve the market segmentation model of K-means clustering, and examines its effectiveness. Further, a personalized recommendation algorithm based on multivariate dynamic user profiles is proposed to recommend products to target users more accurately. A reliable simulation environment is constructed based on the KuaiRec dataset and the classical LastFM dataset to properly evaluate the performance and effectiveness of the model on the recommendation platform. Through the K-means ant colony clustering algorithm proposed in this paper to divide the interest information and attribute information of users, the users as a whole are classified into specific categories, and the online_reward value of the personalized recommendation algorithm based on multivariate dynamic user profiles proposed in this paper fluctuates from 50.05 to 50.49, which is a significantly superior performance. As a result, this paper concludes that cross-cultural marketing strategies should be marketed at four levels: product, price, channel, and promotion, in order to adapt to regional cultures, attract consumers, and build consumer loyalty and satisfaction.

Keywords: k-means clustering, ant colony clustering, cross-cultural, marketing

1. Introduction

With the continuous development and application of Artificial Intelligence (AI) technology, cross-cultural marketing is also facing unprecedented changes, the traditional way of cross-cultural marketing has been unable to meet today's fast-developing market demand, with the help of Artificial Intelligence technology, cross-cultural marketing strategies are being innovated and changed in an unprecedented way [1-4].

 Corresponding author.

E-mail address: Lmylmy54321@163.com (M. Liu).

Received 25 March 2024; Revised 05 June 2024; Accepted 15 December 2024; Published Online 16 April 2025.

DOI: [10.61091/jcmcc127b-299](https://doi.org/10.61091/jcmcc127b-299)

© 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/>).

The application of AI in cross-cultural marketing is mainly manifested in the following methods, firstly, data analysis and prediction, AI can extract useful information from a large amount of data and analyze it to help marketers understand more details about foreign consumers' preferences and behaviors. Through the analysis of big data, it can more accurately predict market trends and consumer demand, so as to optimize product positioning and marketing strategies [5-8]. Secondly, personalized recommendation, AI technology can recommend products and services for individual consumers based on their preferences, purchase history and behavioral patterns. Through the application of machine learning algorithms, AI can conduct in-depth mining of user data and accurately provide each consumer with recommendations that best meet his or her needs, enhancing the user experience and increasing the purchase conversion rate [9-12]. Thirdly, intelligent customer service, the customer service aspect of cross-cultural marketing is crucial to corporate image and consumer satisfaction. AI can realize an intelligent customer service system through technologies such as natural language processing and speech recognition. Such a system can provide round-the-clock and highly efficient customer service to provide consumers with a better shopping and counseling experience [13-16]. At the same time, through the intelligent customer service system to collect and analyze user feedback, enterprises can adjust products and services in a timely manner to improve consumer satisfaction. Fourth, accurate advertising, AI technology can realize accurate advertising through the analysis of user data and network browsing behavior [17-20].

Literature [21] describes the application of AI in marketing and based on a literature review and data clustering using Louvain's algorithm, helps to identify research topics and future research directions aimed at expanding the application of AI in marketing. Literature [22] describes the trajectory of the field of marketing and AI research and develops a context-specific research agenda based on the literature review, describing multiple research paths related to the application and acceptance of AI technologies in marketing, data protection and the role of ethics. Literature [23] examined the use of AI in marketing based on bibliometrics, social network analysis, and other methods, revealing that there are several landmark articles, which constitute the main knowledge streams of AI marketing, and analyzed future directions based on the findings. Literature [24] describes the Artificial Intelligence Marketing (AIM) framework that helps to enhance customer relationships and proposes the strategic use of the AIM framework based on the literature developed with a view to enhance customer relationships including customer trust, satisfaction, and loyalty. Literature [25] analyzes consumer perceptions of AI and AI marketing communications and shows that consumer interpretations of AI are multidimensional and relevant, especially in terms of the functional and emotional aspects of AI, and that AI marketing communications are considered to be unavoidable and acceptable, but that their impact on consumer behavior, among other aspects, is limited. Literature [26] examines AI technology in marketing based on a global perspective and focuses on three dimensions of analysis: country, company and consumer, with a particular focus on the human-computer interaction and automated analytics dimensions of AI technology in marketing, and the interactions between the two dimensions. Literature [27] emphasizes the impact of AI on the field of marketing and examines the application of AI in marketing based on the literature review, illustrates the rising trend of AI in marketing, and indicates avenues for future research in terms of context, methodology and theory. Literature [28] describes the importance of AI for modern marketing, which enables data processing through data analytics, provides insights into consumer behavior and market trends, and enables personalized content and advertisement marketing, effectively improving marketing efficiency, performance, and return on investment. Literature [29] examines the opportunities and challenges of AI in marketing based on the perspective of knowledge creation and knowledge transfer, and analyzes the technological pitfalls and dangers to be aware of when implementing AI, emphasizing that AI will fail to live up to expectations in the marketing field if the challenges it poses are not addressed. Literature [30] describes the positive impacts of AI and emphasizes that its risks to consumers have not yet been taken seriously, based on which it analyzes

the negative impacts of AI on consumers, revealing that these negative impacts mainly include privacy concerns, perceived risks, and other aspects, which provide a reference to enterprise marketing. Literature [31] introduces the relationship between AI and marketing as well as the research related to AI affecting marketing, specifies the attitude of consumers towards AI and the ethical issues brought by AI, and looks forward to the future of AI marketing, and argues the significance of implementing AI marketing to gain a competitive advantage. Literature [32] studied the precision marketing of AI, compared with traditional marketing methods, AI technology achieves a more accurate and personalized marketing effect, and can improve the consumer experience by accurately identifying consumer needs and tapping into potential consumption behavior. The above study analyzes the application of artificial intelligence in marketing, which not only improves the marketing efficiency of enterprises, but also provides consumers with a personalized consumption experience. At the same time, some scholars have pointed out the risks of artificial intelligence in marketing, including privacy and security.

Based on the traditional K-means algorithm, the article proposes an improved K-means ant colony clustering algorithm. The correctness and effectiveness of the improved K-means ant colony clustering algorithm are verified through experiments. Subsequently, based on the market segmentation of the improved K-means ant colony clustering algorithm, the article further proposes a personalized product recommendation algorithm based on multivariate dynamic user profiles. At the feature level, the hidden factor representation and content embedding of users or products are integrated; at the user profile level, the attention mechanism is used to integrate user features and user interests; at the same time, at the user-product matching level, neural networks are used to model more complex relationships to enhance the personalized recommendation capability of the model. Finally, the personalized recommendation algorithm proposed in this paper is subjected to comparison experiments and evaluation experiments.

2. Market segmentation theory and data mining techniques

Transaction data of cross-cultural market users are usually characterized by large data volume and multiple attributes. In this section, after introducing the theory of cross-cultural market segmentation, a K-means ant colony clustering algorithm is proposed for the clustering analysis of such data.

2.1. Definition and steps in cross-cultural market segmentation

Intercultural market segmentation means that an enterprise divides the overall consumer market into several sub-markets with different needs according to different purchasing groups based on consumers' needs for various products. The purpose of cross-cultural market segmentation is to divide the overall market and study a number of categories of sub-markets, to find out the common characteristics of the consumer groups in each sub-market, and then formulate a marketing strategy that is beneficial to the future development of the enterprise, so that the enterprise can find the best way to develop in today's highly competitive market environment.

The general steps of cross-cultural market segmentation are to first analyze and refine the overall market, then study the sub-markets after refinement, select the direction of development in combination with their own actual situation, and finally make the market positioning, the steps of cross-cultural market segmentation are shown in Figure 1. Refinement of the market is the enterprise to the overall market research, and then the enterprise needs to be divided into factors, generalization, digging the latest developments in the existing or future market, which can effectively help enterprises to maximize the understanding of the market [33]. Determine the target market lies in the analysis of the refined sub-markets, so that enterprises according to their own situation to develop more accurate goals, the enterprise's largest resources to play the best results. Market positioning is based on their own situation combined with the development of the identified target submarkets, to develop a

marketing strategy in line with the future development of the enterprise, the fastest to occupy the initiative in the market, so as to avoid wasting a lot of resources to the blind development of the enterprise.

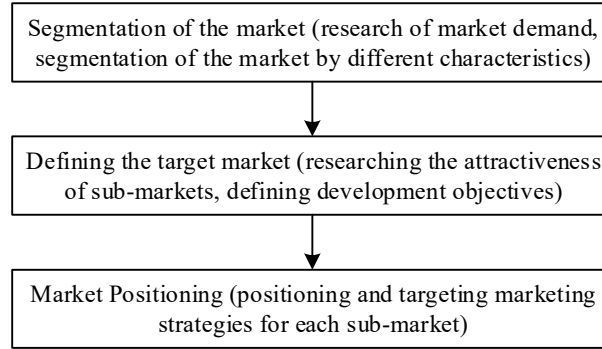


Fig. 1. Market segmentation step

2.2. An improved K-means ant colony clustering algorithm

2.2.1. K-means clustering algorithm. The K-means algorithm is a division-based clustering algorithm that is initialized by giving the number of clusters to be clustered and the clustering centers, and iteratively updating the clustering centers to achieve an optimal solution. The solution is often evaluated using the objective function F :

$$F = \sum_{j=1}^k \sum_{i=1}^m \text{dist}(p_i, c_j) \quad (1)$$

where m is the number of attributes. p_i is the data object belonging to clustering center c_j . dist is the Euclidean distance, defined as follows:

$$\text{dist}_{ij} = \left(\sum_{k=1}^p |x_{ik} - x_{jk}|^2 \right)^{\frac{1}{2}} \quad (2)$$

where p denotes the number of attributes of the data.

K-means The algorithm is described as follows:

- 1) Randomly select k points as the number of clusters to be clustered.
- 2) Randomly select k data objects as initial clustering centers among the data sample points.
- 3) Assign the data objects to the clustering center with the smallest distance according to the distance metric (Euclidean distance is used in the text).
- 4) Recalculate the clustering centers and take the average value of the attributes of the objects in the cluster as the new clustering center.
- 5) Calculate the value of the objective function F using equation (1).
- 6) Repeat steps 3) to 5) until the number of iterations is reached or the clusters stabilize.

When using the *K-means* clustering algorithm one must be given k clusters and an initial cluster centroid at the beginning of the clustering, this k value is very difficult to estimate. Moreover the initial clustering centroid determines an initial division and the subsequent step is to optimize the initial division. The selection of the initial clustering center has a large impact on the clustering results, and once the initial value is not well selected, effective clustering results may not be obtained. In addition, the *K-means* algorithm is very sensitive to isolated data points, and a very small amount of isolated data may have a great impact on the clustering results, thus affecting the performance of the algorithm [34]. However, *K-means* algorithm also has its certain advantages. First of all, its

algorithm idea is simple and easy to implement. In addition, the algorithm has a low computational time complexity and a fast convergence speed. Therefore, the algorithm can be used to preprocess the data and get the prototype of clustering results.

2.2.2. Ant colony clustering algorithm. The idea of the clustering algorithm based on the foraging principle of the ant colony algorithm is as follows: a data sample i is randomly selected as the starting position of the ant traversal, the ant is assigned to the clustering center $c_j (j=1,2,\dots,k)$ according to the formula (3) or randomly generated by a random value less than the pre-set q_0 , then the ants leave pheromone τ_{ij} in the path from the data object i to the clustering center c_j . The probability of ants i choosing the clustering center j is defined as follows:

$$p_{ij} = \frac{\tau_{ij} \times \eta_{ij}}{\sum_{j=1}^k \tau_{ij} \times \eta_{ij}} \quad (3)$$

where η_{ij} denotes the heuristic function, $\eta_{ij} = \frac{1}{d_{ij}}$.

As the ants move, the pheromone on the path is continuously volatilized, and Eq. is defined as follows:

$$\tau_{ij}^{new} = \rho \tau_{ij} + \frac{Q}{d_{ij}} \quad (4)$$

where ρ is the volatilization factor of the pheromone. Q is a positive constant. d_i denotes the distance from data object i to clustering center c_j .

When applying the ant colony clustering algorithm, the speed of convergence of the algorithm is slow, especially at the beginning of the iteration, because the pheromone update is slow so it is difficult to distinguish the pheromone on each path obviously. However, in the late iteration, the pheromone on certain paths keeps piling up, so that the possibility of ants choosing paths for later operations tends more and more to the paths with more pheromone, but there is no guarantee that its solution is the global optimal solution, which leads to the phenomenon of precociousness. Therefore K -means ant colony clustering algorithm is introduced and an improved algorithm is given.

2.2.3. K -means ant colony clustering algorithm. Using the fast convergence property of the K -means algorithm, the cross-cultural market transaction dataset is preprocessed to obtain the clustering centers of the coarse paths. After preprocessing with algorithm K -means, more pheromones are assigned in path d_{ij} (data object to cluster center c_j) according to the cluster center c_j to which data object i belongs τ_y . Next, the ant colony algorithm performs clustering operation by different pheromones on the paths to the respective cluster centers at the beginning. K -means The detailed process of ant colony clustering is described below:

- 1) Randomly select k data from the data sample as the initial clustering center: $c_1, c_2, \dots, c_k$.
- 2) Obtain preprocessed clustering centers and clusters through the K -means clustering process described above.
- 3) Assign initial different pheromone τ_y to each data object i to the corresponding clustering center c_j , pheromone according to the preprocessing process of step clustering (2), the pheromone on each path is different, in which the data samples belonging to a certain class to their clustering centers have relatively high pheromone. Initial Q , ρ (pheromone volatility coefficient), n (number of ants), q_0 (allocation threshold).

4) Randomly generate $p \in (0,1)$. If the value of p is less than the given q_0 , the data samples are assigned to a class according to the principle of pheromone size on the path, and if it is greater than q_0 , the ant transfer probability is calculated according to Equation (3), and the ants are selected to cluster to a class.

5) An ant passes through all data sample points and calculates the value of F , the new clustering center, according to formula (1).

6) All ants complete a traversal to get the better clustering scheme, give the best clustering scheme according to formula (4) for pheromone update.

Reach the predetermined number of iterations then not to move forward, output the optimal solution, otherwise go to step 2) to continue to run.

2.3. Algorithm Validation

In order to facilitate the comparison with the K-means ant colony clustering algorithm, the dataset was first clustered using the traditional K-means algorithm. The parameter K was set to 3, and a total of 100 experiments were conducted, due to the random selection of the initial clustering center, two different divisions appeared in the results of 120 experiments, and the results of the K-means algorithm clustering are shown in Table 1. The results in the table show that the traditional K-means clustering algorithm, in addition to the need to input the expected number of clusters in advance, the improper selection of the initial clustering center can also lead to the clustering results into the local optimum.

Table 1. K-means algorithm clustering results

Partition	Results	Sample size	All kinds of correct Numbers	Accuracy (%)
1	85	625593	516399	82.55
2	35	963641	502447	52.14
Total	120	-	-	82.31

In the improved K-means ant colony clustering algorithm of this paper, the ant colony clustering algorithm can solve the problem of K value determination, and the improved K-means ant colony clustering algorithm can solve the problem of initial clustering center dependence, and at the same time, the initial clustering centers obtained by the ant colony clustering algorithm can accelerate the convergence speed of the subsequent clustering. The sample number of the dataset is 100 and the feature dimension is 4, which is within the processing capability of the ant colony algorithm based on ant colony foraging behavior, so the original dataset is clustered directly using the ant colony clustering algorithm. The ant colony clustering results are shown in Figure 2. The ant colony clustering algorithm clusters the samples in the dataset into 3 classes, and the division of the 3 classes of samples is similar to the original division of the dataset.

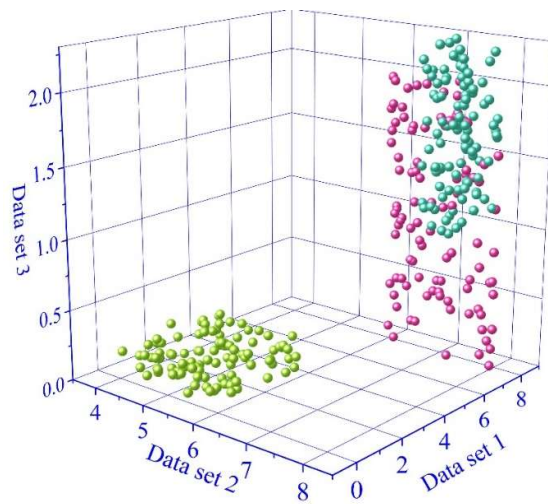


Fig. 2. Ant formation results

Using the number of clusters obtained by the ant colony clustering algorithm3 as the expected number of clusters, a comparison of the initial cluster center selection methods is shown in Table 2. As can be seen from the table, the initial clustering center selection method has almost no effect on the correctness of the results of the improved pheromone-based K-means algorithm, but using the initial clustering centers obtained by the ant colony clustering algorithm can speed up the convergence of the algorithm.

Table 2. The initial cluster center selection method was compared

Initial cluster center source	Random selection(%)	Average time(s)
Ant clustering algorithm output	89.25	0.0091
Random selection	89.25	0.0121

The results of clustering analysis of the dataset using K-means ant colony clustering algorithm are shown in Fig. 3. From the figure, it can be concluded that the K-means ant colony clustering algorithm automatically obtains the appropriate number of class clusters, and the clustering correctness is high and adaptable.

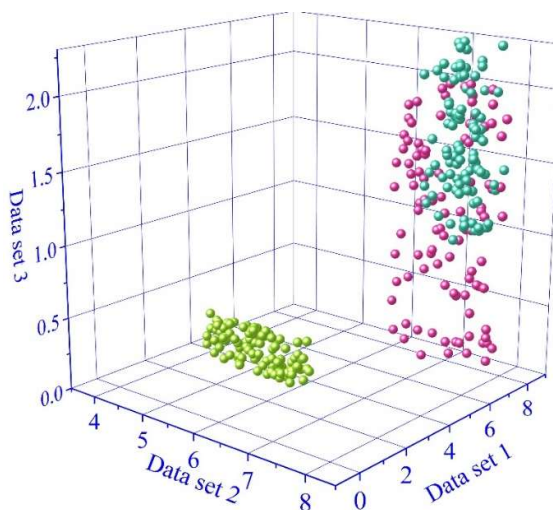


Fig. 3. Cluster analysis results

3. Application of K-means ant colony clustering algorithm in cross-cultural marketing

3.1. Data preparation

This paper utilizes the K-means ant colony clustering algorithm proposed in this paper to carry out practical application research in cross-cultural marketing. The experimental data come from Taobao for a period of time to sell men's wallets store information, sales brand, sales price, the number of reviews, etc., from which the brand sales information table and price sales table were selected as the data for this experiment. Taobao merchandise transaction data set 1 and Taobao merchandise transaction data set 2 were used for the experiment.

3.2. Results of cross-cultural market segmentation

The interest label descriptions are shown in Table 3.

Table 3. Interest label description

Interest label	The meaning of data set 1	The meaning of data set 2
{1,1,1}	They are interested in goods 1, goods 2 and 3	They are interested in goods 1, goods 4 and 5
{0,0,0}	No interest in goods 1, goods 2 and 3	No interest in goods 1, goods 4 and 5
{0,1,0}	Interested in product 2, not interested in product 1 and product 3	Interested in product 4, not interested in goods 1 and goods 5
{1,0,1}	Interested in goods 1 and goods 3, not interested in product 2	Interested in goods 1 and goods 5, are not interested in goods 4
{1,1,0}	Interested in goods 1 and goods 2, are not interested in product 3	Interested in goods 1 and goods 4, are not interested in goods 5
{1,0,0}	Interested in product 1, not interested in product 2 and product 3	Interested in product 1, not interested in goods 4 and goods 5
{0,1,1}	Interested in goods 2 and goods 3, are not interested in product 1	Interested in goods 4 and goods 5, not interested in product 1
{0,0,1}	Interested in product 3, not interested in product 1 and product 2	Interest in goods 5 is not interested in goods 1 and goods 4

The experiment first classifies the users into 8 groups based on their interest information, based on their level of interest in men's wallet goods from 3 different cultures, and obtains the label 1. The user groups are then subdivided based on their attribute information through the K-means clustering algorithm. Taking SSE (clustering evaluation index) as the index, the clustering results under different K values are shown in Fig. 4 (Figs. a~b represent the experimental results of dataset 1 and dataset 2, respectively). It can be found that in dataset 1 and dataset 2, when the K value is greater than 5, the magnitude of SSE change caused by the change of K value becomes obviously smaller, so the experiment selects K=5 to categorize users into 5 groups. In practice, based on the attribute information of users in the same category, the user group characteristics are generalized and summarized to obtain the label 2.



Fig. 4. The clustering results of different k values

As a result of the segmentation of interest and attribute information of the users, combined with label 1 and label 2, the users as a whole were divided into specific categories. The users in each category that are closest to the center are selected as the representative users of the group for subsequent experiments. The distribution of representative users in dataset 1 and dataset 2 is shown in Fig. 5 (Figs. a~b represent the experimental results of dataset 1 and dataset 2, respectively). The visualization of the user data after t-SNE dimensionality reduction can be observed through the figure, in which the rosy red circle in the figure is the representative user, and its distribution in the user group can be seen.

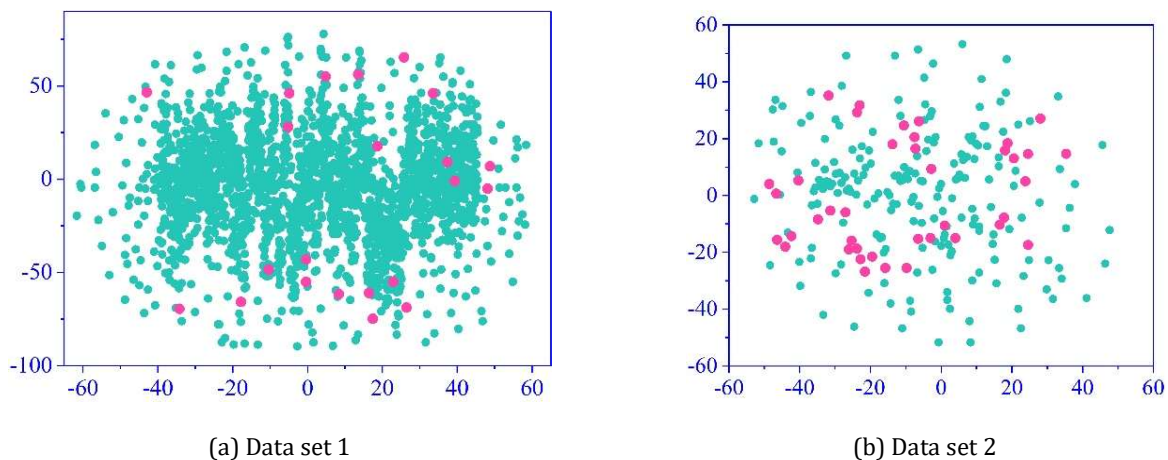


Fig. 5. Data set 1 and number represent the distribution of users in set 2

4. Personalized recommendation algorithms based on multivariate dynamic user profiles

In cross-cultural markets, accurate representation of users and products is crucial. Based on existing research, this section proposes a personalized recommendation algorithm based on multivariate dynamic user profiles to achieve the purpose of recommending products to users in a more diverse and comprehensive way.

4.1. Personalized Recommendation Based on Multiple Dynamic User Profiles

4.1.1. Multiple feature representation module. The principle of the multivariate feature representation module is shown in Fig. 6, which mainly forms a multivariate feature representation based on the fusion of content features and hidden factor features of users or cross-cultural products. In Fig. Q_- , row n of the matrix is the fixed content feature vector extracted from the user or product numbered n after certain feature processing. Row n in the Q matrix is the hidden factor feature vector of the user or product numbered n . The elements in this matrix are initialized to random values and are iteratively updated in the subsequent training process using interaction data, which is essentially a representation extracted from the interaction information [35].

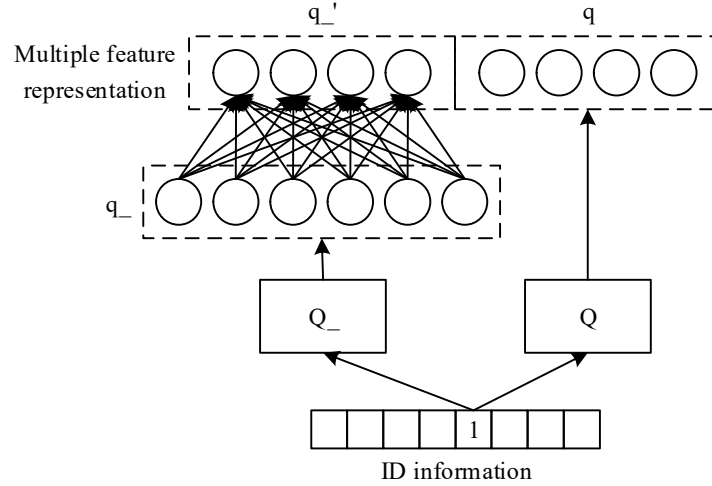


Fig. 6. Multiple characteristics represent module principle block diagram

4.1.2. Multiple dynamic user profile generation module. The principle of the multivariate user profile generation module is shown in Fig. 7, which mainly generates multivariate dynamic user profiles based on the cross-cultural products to be predicted, the collection of historical products, and the multivariate feature representation of users.

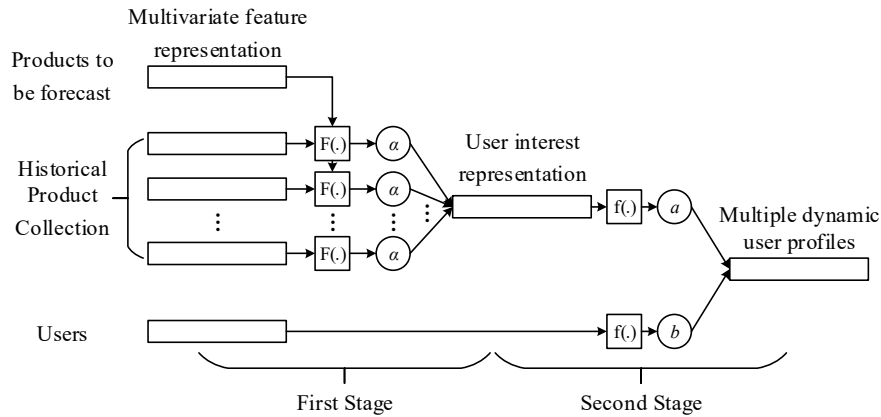


Fig. 7. Multi-user portrait generation module principle block diagram

First, in the first stage, the user interest representation is obtained by adaptive weighted summation of the user's historical product features using the attention mechanism. In the second stage, the user interest representation and user feature vectors are weighted and summed by the attention mechanism to obtain the final multivariate dynamic user representation.

Let the multivariate feature representations of the product to be predicted i , the set of historical products j , and the user k be V_i , I_j , and U_k , respectively, and the user interest representation be H_k , and the multivariate dynamic user representation be M_k . The first stage of the attention mechanism is computed as in Eq. (5), which realizes the process of obtaining the user interest vector H_k based on the adaptive weighting of the user's historical product representation I_j according to the product to be predicted representation V_i .

$$H_k = \frac{1}{|R_k^+|^\alpha} \sum_{j \in R_k^+} a_{ij} I_j$$

$$a_{ij} = \frac{\exp(F(V_i * I_j))}{\left[\sum_{j \in R_k^+} \exp(F(V_i * I_j)) \right]^\beta} \quad (5)$$

where R_k^+ represents the set of products in the user's history, $|R_k^+|$ represents the number of products in the set that need to be weighted and summed, and the larger the value is, the higher the user's activity level is. $*$ denotes the Hadamard product operation, which serves to fuse V_i and I_j . $F(x)$ denotes a forward neural network that contains one hidden layer and has an output dimension of one, and the parameters include W_1 , b_1 , and h_1 , and the specific computational procedure is as in Eq. (6). α is called the interest specification factor, is a real number located in the interval between $[0,1]$, the larger the value, the smaller the modulus of H_k , the lower the calculated click probability, which indicates that the user's activity level influences the user's tendency to choose the product: when $\alpha=1$, the denominator of the formula is equal to the number of summation elements in the formula, the activity level basically does not affect the final prediction of the probability of click. When $\alpha=0$, the denominator is 1, the more active the user is, the higher the predicted probability of a user click. The smaller the value of α , the greater the influence of the user's activity level on the final predicted result. β is called the deflation factor and is located between the intervals $[0,1]$ in order to mitigate the imprecision of computerized numerical representations due to an excessive number of products in the user's history. This chapter mitigates this by narrowing the denominator, i.e., setting $\beta < 1$.

$$F(x) = h_1^T \text{ReLU}(W_1 x + b_1) \quad (6)$$

The second stage of the attention mechanism is computed as in Equation (7), which implements adaptive weighting by determining the reliability of the feature source through the specific representation of the user interest vector U_k and the user features M_k .

$$M_k = a \cdot H_k + b \cdot U_k$$

$$a = \frac{\exp(f(H_k))}{\exp(f(H_k)) + \exp(f(U_k))}$$

$$b = \frac{\exp(f(U_k))}{\exp(f(H_k)) + \exp(f(U_k))} \quad (7)$$

where M_k denotes the final computed multivariate dynamic user portrait, and $f(x)$ denotes a forward neural network that contains one hidden layer and has an output dimension of one, with parameters W_2 , b_2 , and h_2 . The specific computational procedure is shown in Eq. (8).

$$f(x) = h_2^T \text{Tanh}(W_2 x + b_2) \quad (8)$$

4.1.3. User Product Matching Module. The principle of the user-product matching module is shown in Fig. 8, in which the multivariate dynamic user portrait M_k and the multivariate feature representation of the product to be predicted V_i are first spliced to obtain a uniform value vector. The matching score r is then obtained by a multilayer forward neural network to model more complex relationships. In particular, considering the user's personal preference and the product's own quality as well as the total scoring tendency, this model adds bias terms b_i , b_u and b to the neural network output layer.

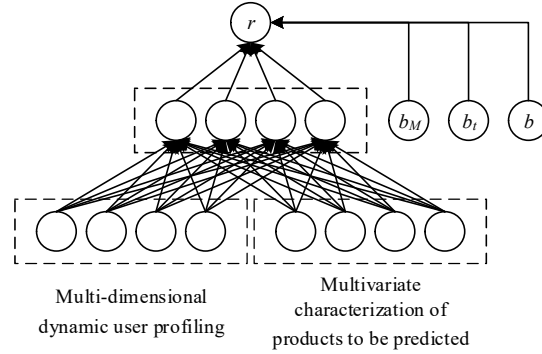


Fig. 8. User product matching module principle diagram

4.1.4. Loss function design. In the recommender system, the parameter optimization problem of the model is equivalent to the problem of minimizing the loss function, in which the commonly used loss functions include the mean square error loss function, the logarithmic loss function, and the Bayesian personalized sorting loss function, which are calculated as follows:

$$Loss_{mse} = \frac{1}{2N} \sum_{i=1}^N (r_i - l_i)^2 \quad (9)$$

$$Loss_{\log} = \frac{1}{2N} \sum_{i=1}^N [l_i \log(r_i) + (1 - l_i) \log(1 - r_i)] \quad (10)$$

$$Loss_{bpr} = -\frac{1}{N} \sum_{i=1}^N \log(\text{sigmoid}(r_{pos} - r_{neg})) \quad (11)$$

where N is the batch size, r_i is the predicted score of the i th test case, l_i is the true score of the i th test case, r_{pos} is the model output corresponding to the positive case, and r_{neg} is the model output corresponding to the negative case. In general, in order to prevent overfitting during training, a L_2 regularization term $\lambda \|\theta\|_2^2$ is usually added to the loss function, where λ is a hyperparameter and θ is a one-dimensional vector containing all parameters in the model.

Among the above loss functions, $Loss_{mse}$ is mainly for the regression problem in the rating prediction task. Most of the recommendation scenarios cannot obtain the exact rating of the user, so it is more suitable for the click prediction task, which belongs to the classification problem in machine learning, with two choices of $Loss_{\log}$ and $Loss_{bpr}$. In this model, from the input of pairs of products to be predicted to the process of calculating r_{pos} and r_{neg} , it involves a large number of network parameter reuse, which greatly affects the readability of the network structure, so the final choice is the logarithmic loss function with a regularization of L_2 , and the specific calculation is as follows:

$$\begin{aligned}
Loss_{\log} = & \frac{1}{2N} \sum_{i=1}^N [l_i \log(r_i) + (1-l_i) \log(1-r_i)] \\
& + \lambda_1 \cdot \|\theta\|_2^2 + \lambda_2 \cdot \|W_1\|_F^2 + \lambda_3 \cdot \|W_2\|_F^2 \\
& + \lambda_4 \cdot (\|P\|_F^2 + \|Q\|_F^2 + \|J\|_F^2)
\end{aligned} \tag{12}$$

where N is the batch size, r_i is the prediction score of the i rd test case, l_i is the true score of the i th test case, θ is the one-dimensional numerical vector containing all the parameters of the forward neural network in the model, W_1 and W_2 are the two-dimensional matrices in the two attention network operations, and P , Q and J are the implicit factor representation matrix of the product to be predicted, the implicit factor representation matrix of the historical product and the user implicit factor representation matrix, respectively. $\|\cdot\|_2$ is the second norm of the vector, and $\|\cdot\|_F$ is the F norm of the matrix.

The optimizer used when the model updates the parameters according to the backpropagation of the loss function shown in Eq. The optimizer used can be the stochastic gradient descent method and its variants. Computationally, Adam is the set of optimizers that takes into account the second-order momentum in Adagrad in addition to the first-order momentum. Therefore, Adam is used as the optimizer for parameter updating in this model.

In terms of network structure, this model unifies the content embedding dimension and the hidden factor vector dimension as embedding size, the dimension of the final multivariate user portrait and the multivariate product features is $2 \times \text{embedding size}$, and the node configuration of the matching network of the user products from the input layer to the output layer is $[4 \times \text{embedding size} \rightarrow 2 \times \text{embedding size} \rightarrow 1]$. In order to prevent the problems of gradient disappearance and gradient explosion during the training process, the activation function of the fully-connected layer is uniformly set to ReLU. In order to match the loss function, the activation function of the output layer is chosen as Sigmoid, which limits the output user product interaction probability to the interval of $(0,1)$. If the content information is missing, the missing vectors can be set to all-zero vectors without changing the network structure.

4.2. Personalized Recommendation Implementation Based on Multiple Dynamic User Profiles

Let the model be $score = match_{Q_u, Q_i, \alpha, \beta, \theta}(user, history, item)$, where $user$, $history$, and $item$ are user ID, user history set ID, and product ID to be predicted, respectively. $score$ is the final matching score.

For parameter initialization, the weight parameters in the forward neural network are initialized with Xavier weights. The bias parameters are initialized to a value of 0 [36]. The parameters in the hidden factor matrix are initialized using random variables obeying $N(0, 0.1)$ a normal distribution truncated by 2σ , i.e., restricted to the interval of $(-0.2, 0.2)$. During the $batch_gen(\cdot)$ process, we adopt a user-centered training format. Referring to the hyperparameter settings in other research works, the initial hyperparameters of this model are set as follows: $embedding\ size = 16$, number of negative samples $K = 4$, learning rate $\eta = 0.0005$, $\alpha = 0$, $\beta = 0.8$.

5. Personalized recommendation algorithm experimental testing

5.1. Experimental design

This section focuses on verifying the effectiveness of the personalized recommendation algorithm based on multivariate dynamic user profiles proposed in this paper in cross-cultural marketing promotion scenarios, as well as the effectiveness of this paper's method on recommendation

performance enhancement. Two cross-cultural marketing datasets, LastFM and KuaiRec, are selected for experiments in this section. The LastFM dataset is a commonly used dataset for product recommendation systems and contains 1,899 users and 17,365 transaction records. KuaiRec is a dataset containing 1,556 users on 3,698 product transaction records jointly released by Taobao and FastFM in 2023. Meanwhile, two machine learning based comparison methods are designed with reference to classical recommendation algorithms. The seven comparison methods include Random, Activefirst, Inactive first, HighRatingfirst, LowRating first, ItemCF, and UserCF. Together, these methods provide rich and reliable control experiments for this study, which help to evaluate the performance of the proposed methods more accurately.

5.2. Experimental results and analysis

5.2.1. Overall experimental comparison. The experimental results of different promotion algorithms are shown in Fig. 9. The figure shows the experimental results of the eight algorithms on three evaluation metrics: product exposure, recommendation algorithm accuracy and recommendation result coverage, where $a(1) \sim a(3)$ and $b(1) \sim b(3)$ correspond to the results on the LastFM and KuaiRec datasets, respectively. Observations lead to the following conclusions:

1) The algorithm in this paper performs much better than the baseline model on the two datasets and on all three evaluation metrics. This indicates that the algorithm in this paper has better comprehensive performance in meeting user needs, improving platform revenue and promoting the overall development of the recommendation platform.

2) In terms of product exposure, the high activity prioritization method and the high rating prioritization method perform relatively well. The poor performance of the low activity prioritization method and low rating prioritization method is mainly attributed to the fact that users with low participation and users with low rating tendency tend to have vague interest preferences, and are therefore less suitable for being the target user groups of promotional activities. The random selection method performs the worst because it does not use any information to optimize the selection strategy.

3) In terms of recommendation algorithm accuracy indicators, UserCF and ItemCF perform relatively well, thanks to the collaborative filtering-based algorithms that fully exploit the similarity information between users and products. Meanwhile, static strategies such as Active first, Inactive first, HighRating first, LowRatingfirst and Random perform relatively poorly.

4) In terms of recommendation result coverage metrics, the High Active first method approach performs well, which is mainly attributed to the frequent interactions between active users and the recommendation platform as well as the rich behavioral data. The interests and preferences of active users are captured more accurately, making them a suitable target user group for promotional activities. In addition, observing the experimental result graphs, it can be found that the result curves of this paper's algorithm show significant fluctuations in all three indicators. Upon analysis, there are two main reasons for this phenomenon: first, the reinforcement learning algorithm needs to strike a balance between exploring unknown states and behaviors and utilizing known information in order to learn the optimal strategy. This leads to instability in the training process. Second, reinforcement learning usually relies on delayed rewards, which makes it difficult for algorithms to directly associate actions with rewards, thus affecting the stability of learning. Nevertheless, in terms of overall trends, the algorithm in this paper achieved superior performance on all three metrics in the later stages of the experiment. This suggests that the algorithm has found a better target user selection strategy in the process of continuous learning and optimization.

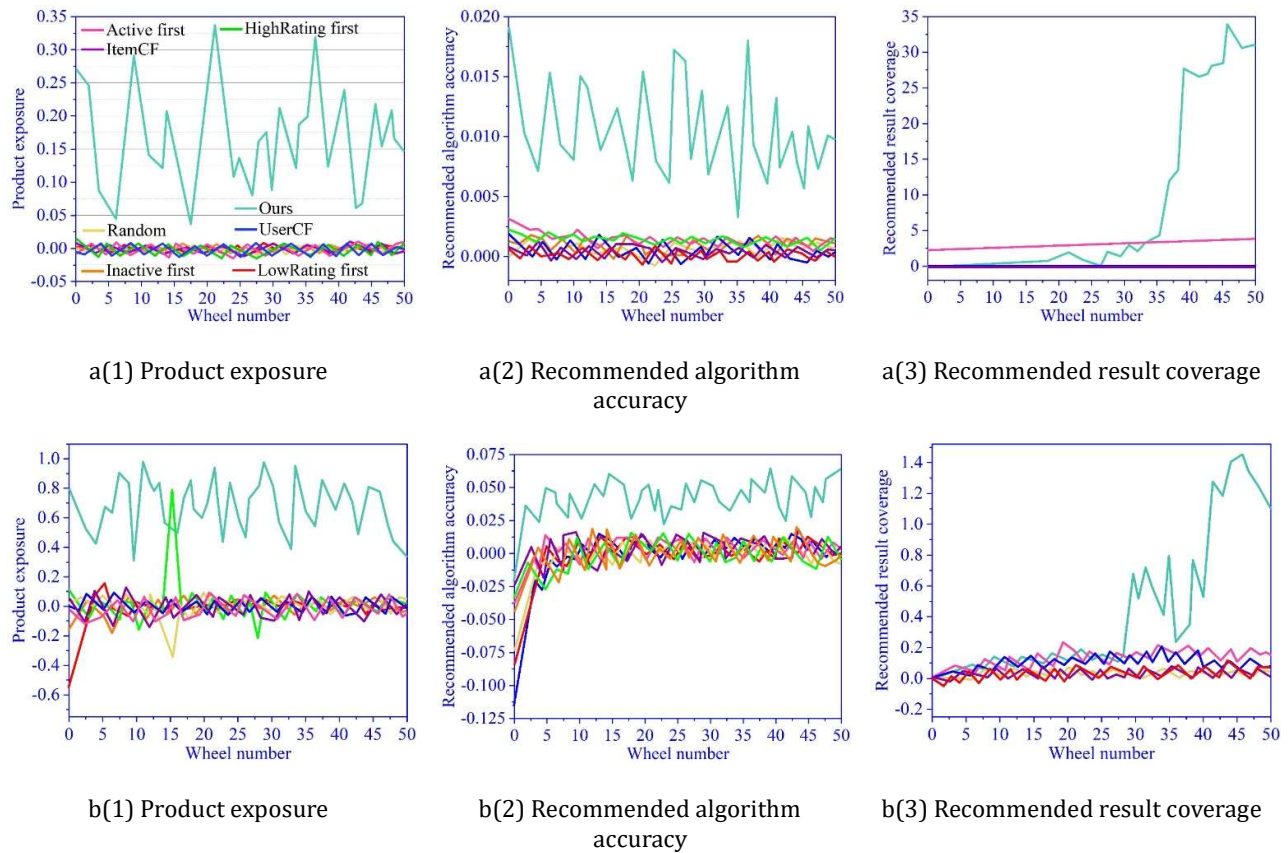


Fig. 9. Experimental results of different extension algorithms

5.2.2. Model evaluation results. A comparison of the evaluation results in the simulation environment is shown in Fig. 10 (Figs. a~b represent the experimental results of this paper's algorithm on LastFM and KuaiRec datasets, respectively). The figure shows the experimental results of this paper's algorithm with seven baseline models on LastFM and KuaiRec datasets. Where the horizontal coordinate represents the number of test rounds and the vertical coordinate represents the user satisfaction of the models in the simulation environment. It is observed that the algorithm of this paper significantly outperforms the baseline model on both datasets. In the LastFM dataset, the performance of the seven baseline models is relatively close to each other, with online_reward fluctuating between 49.55 and 50.03, due to the fact that the strategy adopted in the baseline model cannot adequately capture the dynamic interactions between users and items. The algorithm proposed in this paper performs significantly better, with online_reward fluctuating between 50.05 and 50.49.

In the KuaiRec dataset, the collaborative filtering-based ItemCF performs excellently, thanks to its ability to mine the similarity information between users and items more effectively. Meanwhile, the high rating prioritization method method is significantly better than the low rating prioritization method, which is due to the fact that users who tend to give high ratings to the items are more likely to generate positive feedback, which is conducive to improving user satisfaction. The four algorithms, Active first, Inactive first, Random, and UserCF, perform similarly, with online_reward values fluctuating between fluctuating between 80 and 91. Notably, the model of this paper's algorithm presented in this chapter has an online_reward higher than 140 after a period of training, which significantly outperforms all the baseline models. This can be attributed to the fact that the reinforcement learning based approach is more effective in capturing dynamic environmental changes, as well as in focusing on long-term returns.

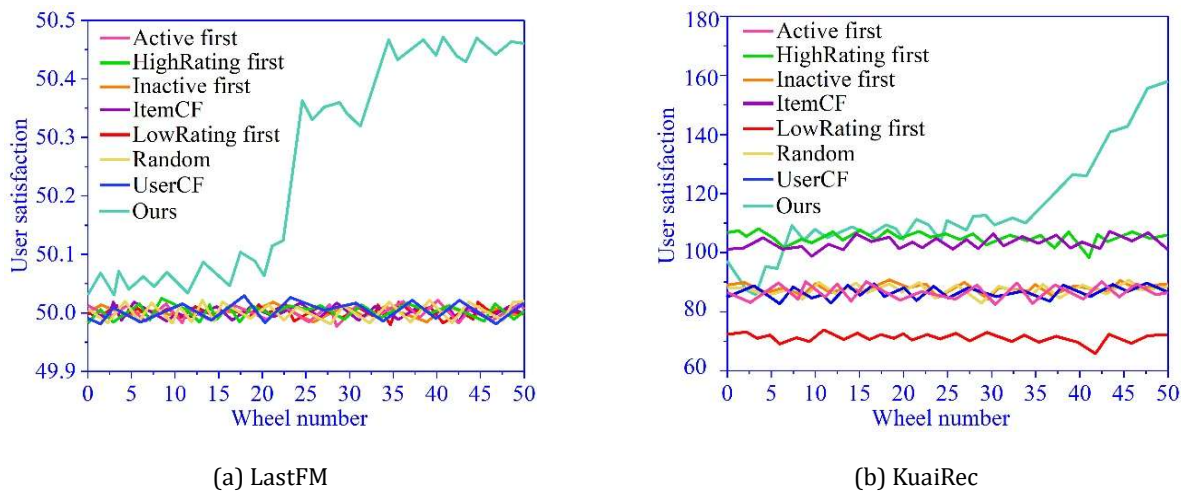


Fig. 10. Comparison of evaluation results in simulation environment

6. Cross-cultural marketing strategies

By using the method of this paper to cluster and divide the cross-cultural market consumer groups, and at the same time to design the personalized product recommendation algorithm for the target users, so that the enterprise can better carry out the marketing strategy planning. This section proposes cross-cultural marketing strategies based on relevant experimental results.

6.1. Product Strategy

Products are the core and foundation of the marketing mix. From a cross-cultural point of view, due to the differences in the cultural environment in which consumers live, and multinational enterprises in China's localized product strategy is designed to produce and sell products from the point of view of the individuality of consumer demand, so in the study and familiarity with the Chinese culture and consumer psychology on the basis of the localized production, product positioning and other aspects of the local market to cater to the needs of the individuality of the thousands of differences, to grasp the degree of difference in the product, to keep the product new and different. The products are new and different. This is conducive to the development of the Chinese market, but also conducive to establish a good international image of the enterprise.

6.2. Price Strategy

Pricing is an important element of the marketing mix, and each culture has its own preferences for pricing strategies as well as the application of methods, with different consumers having different levels of demand for products and cost differences. In the Chinese cultural environment, consumers are more sensitive to price. The price level directly affects the economic efficiency of the enterprise. Facing the changing market demand, when pricing, multinational corporations should not only consider their own cost differences, but also analyze the supply and demand in the Chinese market, consumer income, consumer psychology, and many value tendencies, and consider the impact of non-price factors on consumers' value expectations in the context of the Chinese culture, and adopt corresponding pricing strategies.

6.3. Channel Strategy

From a cultural point of view, the choice of distribution channel is crucial for the consumer. Because the delicate relationship between the distribution model and the consumer is obtained through direct contact, product distribution also plays an important role in cultural penetration. In order for goods

to reach consumers at low delivery costs, TNCs need to focus on building an efficient and cost-effective logistics and distribution network.

6.4. *Promotional Strategies*

The promotional mix is the most culturally conditioned element of the marketing mix and must cater to the local language and culture in the local market. A company's entire promotional mix, i.e., the marketing communication mix, consists of four main tools: advertising, sales promotion, public relations, and personnel promotion. MNCs in China should carefully consider these promotional elements and provide a clear, appropriate and convincing promotional program for the firm and its products as a means of influencing the values of target customers and building relationships with them.

7. Conclusion

The article focuses on the application of K-means and colony clustering algorithm and personalized recommendation algorithm in cross-cultural marketing so that companies can make better marketing strategic planning. The article mainly draws the following conclusions:

- 1) In the cross-cultural market segmentation experiment, the user's interest information and attribute information are divided in data set 1 and data set 2, and the users are divided into specific categories as a whole, the cross-cultural market segmentation is effective, and the experimental results provide good market dynamics for businesses.
- 2) The online_reward value of the algorithm proposed in this paper fluctuates between 50.05 and 50.49, which shows that the personalized recommendation algorithm based on multivariate dynamic user profiles performs better than other models.

References

- [1] Khurram, F. (2023). Role of Artificial Intelligence Bots in Digital Marketing: A Cross-Sectional Study. *Business Review of Digital Revolution*, 3(2), 1-10.
- [2] Chatterjee, S., Chaudhuri, R., & Vrontis, D. (2022). Examining the role of cross-cultural factors in the international market on customer engagement and purchase intention. *Journal of International Management*, 28(3), 100966.
- [3] Singh, N., & Poorvaja, G. (2023). Cross Cultural Integrated Marketing Communication Strategies. Issue 5 *Int'l J.L Mgmt. & Human.*, 6, 1517.
- [4] Liu, Z. (2024). New Dynamics in Cross-Cultural Exchange: A Study on the Role and Impact of Artificial Intelligence Technology in the Global Cultural Market. *Academic Journal of Business & Management*, 6(12), 137-143.
- [5] Odeibat, A. S. (2024). The impacts of artificial intelligence on the future of marketing and customer behaviour. *Cross-Cultural Management Journal*, 26(1), 19.
- [6] Chintalapati, S., & Pandey, S. K. (2022). Artificial intelligence in marketing: A systematic literature review. *International Journal of Market Research*, 64(1), 38-68.
- [7] Huang, M. H., & Rust, R. T. (2021). A strategic framework for artificial intelligence in marketing. *Journal of the academy of marketing science*, 49, 30-50.
- [8] Dimitrieska, S., Stankovska, A., & Efremova, T. (2018). Artificial intelligence and marketing. *Entrepreneurship*, 6(2), 298-304.
- [9] Hermann, E. (2022). Leveraging artificial intelligence in marketing for social good—An ethical perspective. *Journal of Business Ethics*, 179(1), 43-61.

- [10] Mustak, M., Salminen, J., Plé, L., & Wirtz, J. (2021). Artificial intelligence in marketing: Topic modeling, scientometric analysis, and research agenda. *Journal of Business Research*, 124, 389-404.
- [11] Schiessl, D., Dias, H. B. A., & Korelo, J. C. (2022). Artificial intelligence in marketing: A network analysis and future agenda. *Journal of Marketing Analytics*, 10(3), 207-218.
- [12] Wisetsri, W., Ragesh, S. T., Aarthy, C. C. J., Thakur, V., Pandey, D., & Gulati, K. (2021). Systematic analysis and future research directions in artificial intelligence for marketing. *Turkish Journal of Computer and Mathematics Education*, 12(11), 43-55.
- [13] Elhajjar, S., Karam, S., & Borna, S. (2021). Artificial intelligence in marketing education programs. *Marketing Education Review*, 31(1), 2-13.
- [14] Lakshmipriyanka, A., Harihararao, M., Prasanna, M., & Deepika, Y. (2023). A study on artificial intelligence in marketing. *International Journal for Multidisciplinary Research*, 5(3).
- [15] Connell, C., Marciniak, R., & Carey, L. D. (2023). The effect of cross-cultural dimensions on the manifestation of customer engagement behaviors. *Journal of International Marketing*, 31(1), 32-48.
- [16] Manoharan, G., Durai, S., Ashtikar, S. P., & Kumari, N. (2024). Artificial Intelligence in Marketing Applications. In *Artificial Intelligence for Business* (pp. 40-70). Productivity Press.
- [17] Jarek, K., & Mazurek, G. (2019). Marketing and artificial intelligence. *Central European Business Review*, 8(2).
- [18] Senyapar, H. N. D. (2024). Artificial intelligence in marketing communication: A comprehensive exploration of the integration and impact of AI. *Technium Soc. Sci. J.*, 55, 64.
- [19] Vishnoi, S. K., Bagga, T. E. E. N. A., Sharma, A. A. R. U. S. H. I., & Wani, S. N. (2018). Artificial intelligence enabled marketing solutions: A review. *Indian Journal of Economics & Business*, 17(4), 167-177.
- [20] Noranee, S., & bin Othman, A. K. (2023). Understanding consumer sentiments: Exploring the role of artificial intelligence in marketing. *JMM17: Jurnal Ilmu ekonomi dan manajemen*, 10(1), 15-23.
- [21] Verma, S., Sharma, R., Deb, S., & Maitra, D. (2021). Artificial intelligence in marketing: Systematic review and future research direction. *International Journal of Information Management Data Insights*, 1(1), 100002.
- [22] Vlačić, B., Corbo, L., e Silva, S. C., & Dabić, M. (2021). The evolving role of artificial intelligence in marketing: A review and research agenda. *Journal of business research*, 128, 187-203.
- [23] Chen, J., Ablanedo-Rosas, J. H., Frankwick, G. L., & Arévalo, F. R. J. (2021). The state of artificial intelligence in marketing with directions for future research. *International Journal of Business Intelligence Research (IJBIR)*, 12(2), 1-26.
- [24] Yau, K. L. A., Saad, N. M., & Chong, Y. W. (2021). Artificial intelligence marketing (AIM) for enhancing customer relationships. *Applied Sciences*, 11(18), 8562.
- [25] Chen, H., Chan-Olmsted, S., Kim, J., & Mayor Sanabria, I. (2022). Consumers' perception on artificial intelligence applications in marketing communication. *Qualitative Market Research: An International Journal*, 25(1), 125-142.
- [26] Kopalle, P. K., Gangwar, M., Kaplan, A., Ramachandran, D., Reinartz, W., & Rindfleisch, A. (2022). Examining artificial intelligence (AI) technologies in marketing via a global lens: Current trends and future research opportunities. *International Journal of Research in Marketing*, 39(2), 522-540.
- [27] Labib, E. (2024). Artificial intelligence in marketing: exploring current and future trends. *Cogent Business & Management*, 11(1), 2348728.

- [28] Umamaheswari, D. D. (2024). Role of Artificial Intelligence in Marketing Strategies and Performance. *Migration Letters*, 21(S4), 1589-1599.
- [29] De Bruyn, A., Viswanathan, V., Beh, Y. S., Brock, J. K. U., & Von Wangenheim, F. (2020). Artificial intelligence and marketing: Pitfalls and opportunities. *Journal of Interactive Marketing*, 51(1), 91-105.
- [30] Barari, M., Casper Ferm, L. E., Quach, S., Thaichon, P., & Ngo, L. (2024). The dark side of artificial intelligence in marketing: meta-analytics review. *Marketing Intelligence & Planning*, 42(7), 1234-1256.
- [31] Makhlooq, A., & Al Mubarak, M. (2024). Artificial intelligence and marketing: Challenges and opportunities. *Technological Innovations for Business, Education and Sustainability*, 3-16.
- [32] Yang, X., Li, H., Ni, L., & Li, T. (2021). Application of artificial intelligence in precision marketing. *Journal of Organizational and End User Computing (JOEUC)*, 33(4), 209-219.
- [33] Jiajie Liu, Tao Liu & Guangzhong Cao. (2025). Examining Rental Housing Market Segmentation in Chinese Megacities: The Case of Beijing. *Applied Spatial Analysis and Policy*, 18(1), 31-31.
- [34] Yu Kang & Qiang Wu. (2020). Application of ant colony clustering algorithm in coal mine gas accident analysis under the background of big data research. *Journal of Intelligent & Fuzzy Systems*, 38(2), 1381-1390.
- [35] Shuping Zhang. (2025). Integrating User Profiles and Collaborative Filtering for Smart Recommendation of Tourism City Cultural and Creative Products. *International Journal of High Speed Electronics and Systems*, (prepublish).
- [36] Ju Zhou & Hong Li. (2024). Multi-dimensional Model of Python Resources Based on Portrait Technology: Research on Building and Optimizing Recommendation Services. *The Frontiers of Society, Science and Technology*, 6(3).