



Journal of Advanced Research in Computing and Applications

Journal homepage:
<https://karyailham.com.my/index.php/arca/index>
ISSN: 2462-1927



Optimizing ChatGPT Interactions through Understanding User Preferences and Emotions by Machine Learning and User Profiling

Liang Jiayi¹, Maslina Zolkepli^{1,*}, Siti Nurulain Mohd Rum¹

¹ Faculty of Computer Science and Information Technology, Universiti Putra Malaysia, Selangor, Serdangor, 43400 UPM, Malaysia

ARTICLE INFO

Article history:

Received 2 August 2025
Received in revised form 8 August 2025
Accepted 17 September 2025
Available online 12 October 2025

Keywords:

ChatGPT; User emotions; Word Cloud;
User Profiling; Random Forest; XGBoost;
Multinomial Naive Bayes; Logistic
Regression; K-means

ABSTRACT

The exploration of the unknown by people has spurred ChatGPT's evolution. However, since users' expressions are usually vague and inaccurate, ChatGPT also overlooks the heterogeneity of users, resulting in users being dissatisfied with ChatGPT's responses. Counting dissatisfied user segments and increasing satisfaction involves employing machine learning and User Profiling. The design and implementation of the whole experiment mainly include the collection, storage and processing of users and review data. using Word Cloud and machine learning algorithms for feature selection and analysis to obtain feature importance. In addition to applying the clustering method to identify different user groups of ChatGPT. Among them, dynamic real-time data collection is mainly done using the distributed message queue Kafka, persistent storage of data is achieved by the log collection tool Flume, and data processing using real-time computation and low-latency streaming computation is mainly done by the distributed computing engine Flink. It turned out that the term user experience was the most important feature to improve user satisfaction, and that the target group of ChatGPT that demanded to improve user satisfaction was the group of users with high demand and low satisfaction. These steps are validated to be essential for improving response speed, aiding prompt engineers to deal with issues, and ensuring the sustained growth and maximum benefits of ChatGPT.

1. Introduction

People rely on information daily, from weather updates and scientific facts to everyday decisions like choosing clothes or food. ChatGPT, a sophisticated language model, has emerged as a valuable tool to assist in these information needs. Trained on a massive dataset of text and code, ChatGPT can generate human-quality text, translate languages, and answer questions in an informative way proposed by Bansal, R [1].

While ChatGPT is a powerful tool, it's important to recognize its limitations. It may sometimes generate incorrect or misleading information, especially when presented with biased or inaccurate data proposed by Bansal, R [1]. Additionally, ChatGPT may struggle to fully understand and respond to nuanced queries, as it often selects responses based on statistical probabilities rather than deep

* Corresponding author.

E-mail address: masz@upm.edu.my

<https://doi.org/10.37934/arca.40.1.117127>

contextual understanding. To improve the user experience, future developments should focus on enhancing ChatGPT's ability to comprehend user intent and tailor responses accordingly.

As ChatGPT continues to evolve, a new profession has recently been born called Prompt Engineer (PME), which allows people to better meet the needs of users and targeting help to those with high-need users. Prompt Engineering to optimize the interaction with ChatGPT to improve its efficiency and output quality, proposed by Mishra, D. *et al.*, [8]. The Prompt Engineer needs to carefully craft questions to ensure they are both specific and clear so that ChatGPT can provide accurate and relevant answers.

ChatGPT often fails to account for individual user differences proposed by Iyer, A. A., & Vojjala, S. [2]. This means that even when users express similar needs, their desired outcomes can vary significantly. The model's inability to fully grasp the unique contexts and perspectives of each user limits its ability to provide truly personalized responses. For instance, two users seeking travel advice may have distinct preferences and constraints, yet ChatGPT might offer a generic response that fails to address their specific requirements.

On the other hand, there has been focus on the gap between users' expressions and their underlying needs proposed by Phang, J. *et al.*, [3]. Users often have a clear idea of what need but may not articulate the queries precisely. This vagueness leads to challenges for ChatGPT, which relies heavily on the clarity and specificity of user inputs to generate relevant and accurate responses. Therefore, when users express the needs ambiguously, ChatGPT struggles to infer and address the implicit aspects of the problem.

Additionally, the uniformity of responses provided by ChatGPT, despite the flexibility in terms of accessibility and usage, further complicates the user experience. Users are often presented with generic answers and are left to sift through these to find the information that specifically applies to a unique situation. This one-size-fits-all approach does not align well with the diverse and personalized needs of users, making it imperative for future developments in AI communication models to focus more on personalization and contextual understanding.

The study will investigate two primary questions: first, which specific feature of ChatGPT is most crucial in enhancing user satisfaction, as determined by user sentiment feedback. Second, it will identify the target demographic group that would benefit most from improvements to ChatGPT's user satisfaction. To achieve these objectives, the study will compare the accuracy of various machine learning models in analyzing user emotional feedback on ChatGPT. Additionally, it will employ clustering methods to identify distinct user segments of ChatGPT.

The development of a sentiment analysis model for ChatGPT marks a significant advancement in the fields of Natural Language Understanding (NLU) and sentiment computing, proposed by Bansal, R [1]. This model tackles the intricate challenge of discerning emotional context in conversational AI interactions, where identifying sentiment can be particularly difficult due to the conversational and often ambiguous nature of dialogue. By enhancing real-time understanding of user emotions, the model promises to facilitate more empathetic and effective human-computer interactions. This breakthrough holds substantial potential for enhancing user experiences across various domains, including customer service, mental health applications, and personalized content delivery.

The introduction of User Profiling addresses the challenge of filtering relevant information from a plethora of responses proposed by Khan, Z. A. *et al.*, [4]. By analyzing User Profiling, the system can accurately identify high-need users. This targeted approach not only saves users time in selecting relevant information but also enables precise personalized content delivery.

In summary, current research is limited by the following: First, the ambiguity of user expressions and the limitations of ChatGPT. ChatGPT's generalized responses cannot meet users' unique contextual and personalized needs. Second, there is a lack of in-depth understanding of user

emotions. Third, existing research has underutilized user data. Therefore, this study aims to address these gaps. The main goals are first to identify the key features that influence user satisfaction and second to identify target user groups. The findings of this study are of great significance for improving ChatGPT's performance and user experience, which can enhance user satisfaction, target improvements and support, and thus improve the user experience. They provide a foundation for the development of ChatGPT personalization, which is crucial for developing more personalized and context-aware AI communication models that can bridge the gap between user expressions and their underlying needs. This not only provides guidance for the direction of technological development but also prompts attention to the ethical issues of AI in emotion understanding and human-computer interaction.

CustomerID	Gender	Age	Income	Score
1	Male	19	15	39
2	Male	21	15	81
3	Female	20	16	6
4	Female	23	16	77
5	Female	31	17	40
6	Female	22	17	76
7	Female	35	18	6
8	Female	23	18	94
9	Male	64	19	3
10	Female	30	19	72
11	Male	67	19	14
12	Female	35	19	99
13	Female	58	20	15
14	Female	24	20	77
15	Male	37	20	13
16	Male	22	20	79
17	Female	35	21	35
18	Male	20	21	66
19	Male	52	23	29
20	Female	35	23	98
21	Male	35	24	35
22	Male	25	24	73
23	Female	46	25	5

Fig. 1. Sample dataset of ChatGPT_Customers

date	title	review	rating
2023-05-21 16:42:24	Much more accessible for blind users than ..	Up to this point I've mostly been using Ch..	4
2023-07-11 12:24:19	Much anticipated, wasn't let down.	I've been a user since it's initial roll o..	4
2023-05-19 10:16:22	Almost 5 stars, but.. no search function	This app would almost be perfect if it was..	4
2023-05-27 21:57:27	4.5 stars, here's why	I recently downloaded the app and overall..	4
2023-06-09 07:49:36	Good, but Siri support would take it to th..	I appreciate the devs implementing Siri su..	4
2023-05-31 10:20:48	App review	No doubt, this technology is absolutely i..	1
2023-06-23 08:10:54	Almost perfect except for..	Please provide a TABLET experience on iPad..	3
2023-07-04 19:02:08	Chat GPT: The Underrated Buddy in Your Poc..	Chat GPT is seriously underrated, dude! I ..	4
2023-05-18 21:10:34	Nice and quick!	On this app, as opposed to on the website..	5
2023-06-15 15:39:10	The app of all apps for AI	There's been times of apps tooting their ar..	5
2023-06-08 16:49:52	A no-brainer	This is the beginning of the future of how..	5
2023-06-13 15:25:35	Grateful this AI isn't an actual Hologram ..	I couldn't resist and finally downloaded t..	5
2023-05-25 14:25:42	Can't use on iPad	I installed this on my iPhone and it works..	2
2023-05-18 17:57:57	OpenAI iOS App Review	The OpenAI iOS App is a fantastic tool for..	5
2023-05-18 23:43:02	It's a great start but please add shortcut..	Works well, and I enjoy the haptic feedbac..	4
2023-05-22 00:37:56	Wow! What a gift. I rarely write reviews. ..	Dear OpenAI Team, ☺️ • I would like to sha..	5
2023-07-24 21:22:03	Wow! What a gift. I rarely write reviews. ..	I very rarely give reviews but want to wit..	5
2023-05-18 16:58:37	A Leap into the Future with the ChatGPT 10..	The ChatGPT iOS App is an ingenious innova..	5
2023-05-18 21:20:26	Exceptional, Unbeatable App Experience!	In a digital landscape teeming with applic..	5
2023-05-18 21:52:14	The other sources weren't giving responsib..	I act as a Learning coach for my child who..	5
2023-05-25 06:54:38	A step forward	I like the idea of having an app as oppose..	4
2023-06-24 13:54:59	ChatGPT is impressive	ChatGPT is an impressive AI language model..	5

Fig. 2. Sample dataset of ChatGPT_reviews

2. Methodology

2.1 Dataset

The study utilized MySQL database to store 200 pieces of user data, including user number, gender, age, investment in ChatGPT (in dollars), and ChatGPT scores (percentile scale) as in Figure 1. Analysis of this data reveals user demographic characteristics, investment behavior, and satisfaction with ChatGPT. Despite the small amount of data, this information is still valuable in understanding user behavior and optimizing product strategy. The variable "Investment in ChatGPT (Income)" is represented user needs in US dollars, this variable shows the amount invested by users in ChatGPT. The range being from 0 to infinity indicates a wide spectrum of user engagement and financial commitment. ChatGPT Score on a hundred-point scale, based on the data collected, these range from 3 to 99, this metric reflects users' ratings, satisfaction or evaluations of ChatGPT. Given the diverse nature of these data points, analyzing this dataset can provide valuable insights into engagement levels and satisfaction with ChatGPT. The relatively small dataset size (200 entries) may limit the statistical significance of the findings. However, it can still provide valuable preliminary insights

Storing the data into MySQL database which has four columns including date, title, review, and rating, a total of 2670 data was collected from the aforementioned 200 users as in Figure 2. The date is the time when the data was uploaded, which was distributed from March to July 2023. The title is the subject of the upload, which summarizes the content of the review. Review column is a rich repository of qualitative data, comprising various formats like text, images, and diverse writing styles, including uncase sensitivity and irregular formats, reflecting the varied and authentic voices of customers. The rating is the user's rating, with 1 being the lowest and 5 being the highest. The primary

objective of data preprocessing is to handle titles and comments, which involves removing irrelevant symbols and expressions, cleansing the data, and enhancing the accuracy of sentiment analysis. Filtering out common words that contribute little to sentiment analysis helps focus on more meaningful content. Simplifying words into their root forms standardizes the dataset and consolidates word variations. Carefully eliminating sensitive information maintains privacy and ethical standards. This structured and thorough approach to data storage and preprocessing is crucial for the success of research. It ensures data quality and consistency, laying a solid foundation for advanced data analysis techniques such as sentiment analysis.

2.2 Sentiment Analysis

Based on the flowchart, the experiment can be divided into two steps. The first step is data processing, and the second step is creating word clouds, validating word clouds, and obtaining feature importance. The flowchart for sentiment analysis is shown in Figure 3.

The initial phase of data processing relies heavily on the Linux system. To establish the necessary environment, the Aliyun server is installed and configured with a 64-bit Linux system proposed by Liu [15]. Subsequently, Flume is installed and configured to capture and send files to Kafka proposed by Lee. et al., [16]. Flink then consumes data from Kafka, processes it, and stores the results in MySQL[17,27]. The data stored in MySQL will undergo processing, also known as feature engineering proposed by Esh, M [5]. This includes data cleaning, such as standardizing letter cases and removing special symbols and expressions, stopwords, and sensitive vocabulary. Feature selection, such as choose sentiment scoring and comment to analysis. Feature transformation, such as Word2vec word vector model, etc.

Word Cloud is generated as part of the data analysis. This visual representation highlights the most frequent words within the reviews, which can provide insights into common themes or sentiments introduced by Kim, M. [6]. Then, four machine learning algorithms, labeled as Multinomial Naive Bayes, Logistic Regression, Random Forest, and XGBoost, will be employed to analyze the reliability of the word clouds. By comparing the performance metrics of different algorithms, the best model will be determined, thus obtaining the feature importance of user reviews.

Performance metrics quantify a model's effectiveness in machine learning and statistical analysis, aiding in comparing different models or algorithms. They vary based on the model and analysis goals. Common metrics in sentiment analysis include Accuracy, Precision, Recall, and F1-Score introduced by Rawat, P. et al., [7]. Accuracy measures overall correctness, while Precision assesses classification accuracy for a specific sentiment. Recall gauges a model's ability to capture a specific emotion category, and F1-Score combines Precision and Recall to provide a balanced measure of overall performance, especially useful with uneven sample distributions. These metrics offer unique insights into model performance, customizable to suit specific objectives and dataset characteristics. It's crucial to select metrics aligned with model goals and data nature.

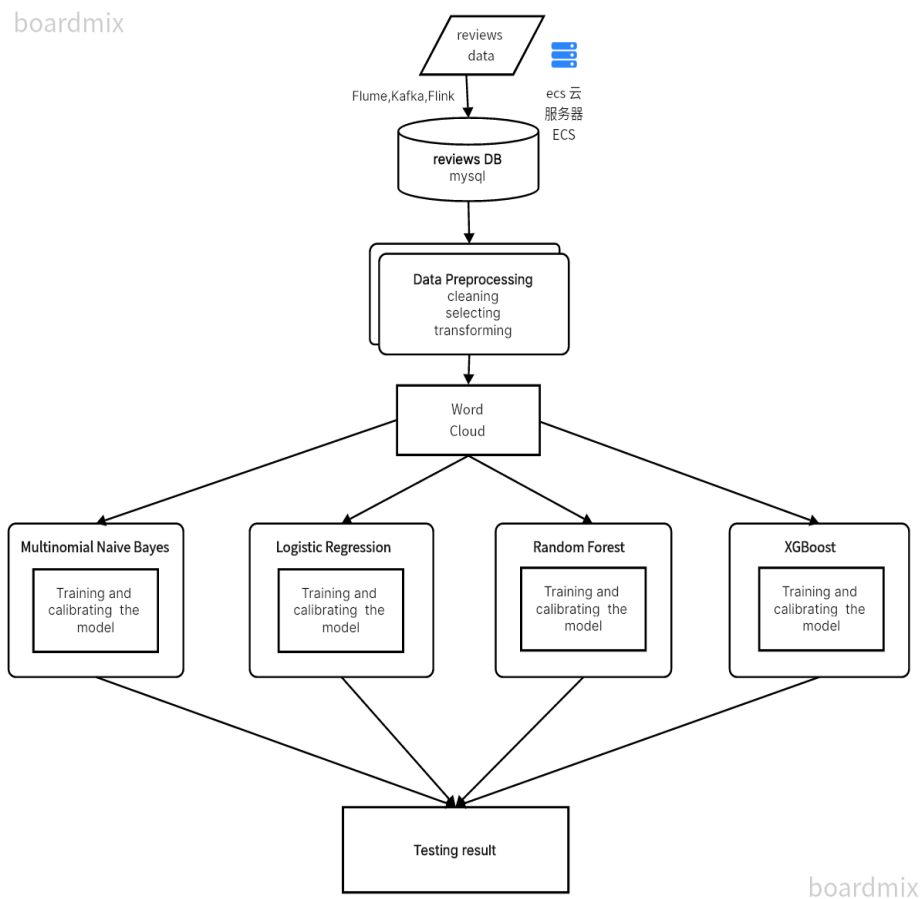


Fig. 3. The flowchart for sentiment analysis

2.3 User Profiling

User Profiling is labeled user models abstracted from users' social attributes, habits, and consumption behaviors. It uses the concept of hierarchical structure (includes user profiles, dimensions, and labels) to present the distribution of features of crowd data in a clear and organized way. The core work of building User Profiling is labeling users. The steps to build User Profiling usually include, data collection, data cleaning, data standardization, user modeling, tag mining and tag verification introduced by Khan *et al.*, [4].

Since the acquired data is relatively simple, extensive processing is not required, and it can be used directly. With the data prepared, user modeling is performed using the K-means algorithm. This unsupervised machine learning model is particularly effective for segmenting users into distinct groups based on their attributes and behaviors. Based on the clustering results from the K-means model, users are classified into different categories. This involves understanding the characteristics and preferences of each user group, which can be vital for personalized marketing, product development, or enhancing user experience.

Table 1
Machine Learning Model Performance Comparison

	Multinomial Naïve Bayes			Logistic Regression			Random Forest			XGBoost		
	Precision	Recall	F1-score	Precision	Recall	F1-score	Precision	Recall	F1-score	Precision	Recall	F1-score
Negative	0.98	0.61	0.73	0.78	0.87	0.82	0.84	0.92	0.88	0.92	0.90	0.91
Neutral	0.57	0.73	0.64	0.90	0.76	0.83	0.93	0.78	0.85	0.92	0.93	0.92
Positive	0.66	0.89	0.76	0.82	0.90	0.86	0.80	0.90	0.85	0.90	0.90	0.90
Accuracy	0.71			0.83			0.86			0.91		

As shown in Table 1, Multinomial Naive Bayes was driven by its suitability for processing text data, probabilistic nature for classification tasks, and simplicity in implementation introduced by Graham, C. *et al.*, [10]. The Multinomial Naive Bayes model yielded an overall accuracy of 71%, indicating reasonable performance across the test set. It exhibited a high recall for positive sentiment (0.89) but a lower recall for negative sentiment (0.61), suggesting potential improvements in identifying negative sentiment data. Logistic Regression was chosen for its interpretability and ability to serve as a baseline model.

Logistic Regression achieved an overall accuracy of 83%, with high precision for the neutral category (0.90) and high recall for the positive category (0.90), demonstrating balanced performance across categories. But the performance in negative sentiment the precision, recall and F1--score all lowest. So maybe the Word Cloud of sentiment does not linear relationship. Random Forest was selected for its adaptability to different data distributions and ensemble learning capabilities. Random Forest achieved an accuracy of 86%, excelling in precision for the neutral category (0.93) and recall for the negative category (0.92). In order to get higher accuracy, XGBoost was chosen for its high accuracy, optimization capabilities, and ease of implementation. XGBoost demonstrated balanced and efficient performance with precision, recall, and F1-scores above 90% for each category, achieving an accuracy of 91% across the test set.

Overall, each model demonstrated strengths in different aspects of performance, with XGBoost exhibiting the highest overall accuracy and balanced performance across categories.

3.2 User Profiling

3.2.1 Relationship Analysis

Analyzing user data on age, income needs, satisfaction scores, and gender reveals several key insights. Firstly, there is a consistent demand for humanization across all age groups, suggesting a universal desire for ChatGPT to exhibit more human-like qualities in its interactions. Secondly, the majority of user satisfaction scores fall within the average range, indicating a lack of extreme positive or negative experiences. Thirdly, the user base primarily consists of individuals between 20 and 70 years old, with a concentration on younger and middle-aged demographics. Interestingly, age does not appear to significantly influence demand for ChatGPT, suggesting its appeal to a broad age range. Finally, a potential gender disparity emerges, with men demonstrating a higher demand for ChatGPT compared to women. This suggests a greater need for user-friendliness and high performance among male users.

3.2.2 User Modeling

Figure 7 presents an elbow diagram, a common technique for determining the optimal number of clusters in K-means clustering introduced by Li *et al.*, [14]. The diagram illustrates the relationship between the number of clusters and the inertia value, which measures the sum of squared distances between data points and their assigned cluster centers. As the number of clusters increases from 2 to approximately 10, the inertia value decreases rapidly, indicating tighter and more compact clusters. However, beyond this point, the rate of decrease slows down, signifying diminishing returns. The optimal number of clusters is typically identified at the "elbow point," where the curve begins to flatten. In this analysis, the elbow point is determined to be 5, suggesting that five clusters provide an appropriate balance between cluster compactness and overall model complexity.

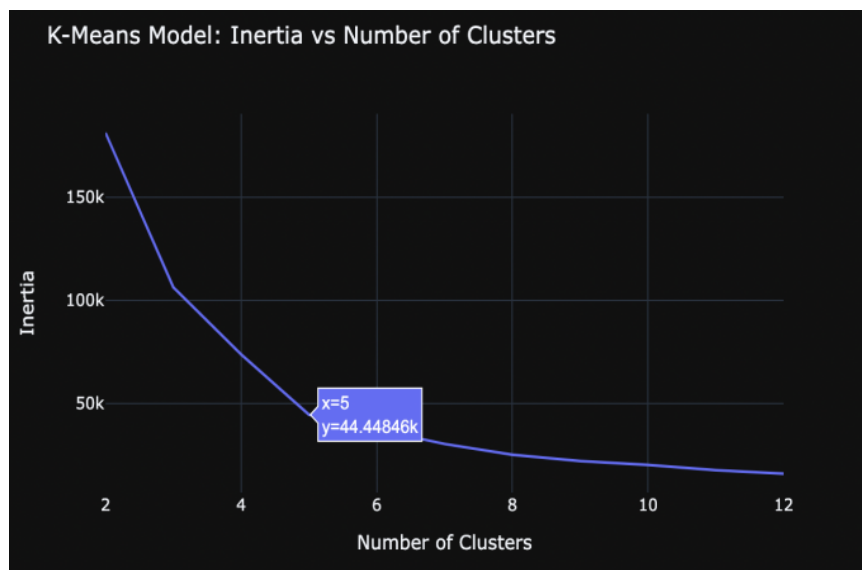


Fig. 7. User modeling with K-means.

3.2.3 Label Mining

The code outputs clustering labels (integers 0 to 4) for each data point, representing assigned groups. The resulting array displays these labels for each data point. For instance, as shown in Figure 8, the first data point is in cluster 4, the second in cluster 2, and so forth. This reveals which cluster each data point is assigned based on characteristics introduced by Samih *et al.*, [23]. This insight enables analysis of group characteristics, creation of User Profiling, and development of tailored strategies or services. Data points in the same cluster represent similar demographic customers, facilitating targeted campaigns.

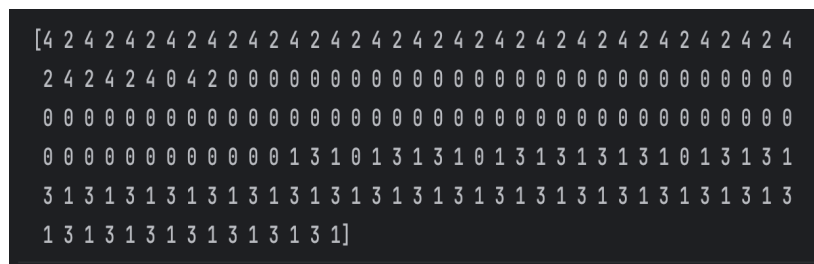


Fig. 8. The classification of clustering labels.

Use users' investment in ChatGPT to represent users' demand for ChatGPT. In the study, 81 individuals, labelled as 0, indicate that increased investment in ChatGPT does not necessarily result in a better user experience, for example, unmet customer needs even after upgrading to ChatGPT version 4.0. With 39 individuals labelled as 1, the second largest group, users rate their ChatGPT experience higher. This indicates that ChatGPT's responses can satisfy the needs of such users, potentially due to broad questions not requiring precise answers or users utilizing effective questioning methods. Tags 0 and 1 collectively represent individuals with high ChatGPT expectations, comprising 60% of the study population. This underscores the growing demand for a more empathetic artificial intelligence like ChatGPT, signaling a necessary direction for future development introduced by Aziz *et al.*, [22].

Table 2
The result of clustering labels

Label Number	Number of People	Representative Group
0	81	High Input, Low Score
1	39	High Input, High Score
2	22	Medium Input, Medium Score
	35	Low Input, Low Score
4	23	Low Input, High Score

Comparing the 35 individuals who rated ChatGPT as 3 (less humane) with the 23 who rated it as 4 (more humane) highlights a low sensitivity to less empathetic versions. This further underscores the strong user preference for a more empathetic ChatGPT, as evident in Table 2.

4. Conclusions

Analyzing emotional responses to ChatGPT revealed widespread negative sentiment towards "user experience" in a Word Cloud. Four machine learning algorithms were employed to assess Word Cloud accuracy and analyze feature importance. To address the identified issues, can recommend categorizing users based on the specific needs. Findings indicate dissatisfaction among high-demand users with ChatGPT's performance, suggesting additional support such as prompt engineering training to enhance user experience. Data analysis confirmed negative sentiment towards user experience, while user profiling identified specific needs of high-demand users, showcasing dissatisfaction with ChatGPT's current performance.

In terms of data, the demand is to seek cooperation with ChatGPT to gain access to real-time dynamic data, enabling comprehensive testing and refinement of the sentiment analysis module. And explore alternative data collection methods (e.g., APIs, crowd-sourcing proposed by Flores, H. I. O. [9]) to supplement existing user information and achieve more robust user modeling. The long-term goal is to conduct longitudinal studies with larger datasets to validate the model's effectiveness and generalizability.

As for the model itself, it needs to be combined with advanced NLP techniques (e.g., contextual analysis, sarcasm detection) to increase the accuracy and nuance of sentiment analysis. Integrate the sentiment analysis module with downstream applications to demonstrate real-world impact (e.g., customer feedback analysis, and social media monitoring). At the same time, investigate the optimal methods for Prompt engineering and user profiling within the ChatGPT context.

Acknowledgement

This research was supported by the Faculty of Computer Science and Information Technology, Universiti Putra Malaysia.

References

- [1] Bansal, R. *Unveiling the Potential of ChatGPT for Enhancing Customer Engagement*. In *Leveraging ChatGPT and Artificial Intelligence for Effective Customer Engagement*, 111–128. IGI Global Scientific Publishing, 2024.
- [2] Iyer, A. A., and S. Vojjala. "Augmenting Sentiments into Chat-GPT Using Facial Emotion Recognition." In *2024 10th International Conference on Advanced Computing and Communication Systems (ICACCS)*, vol. 1, 69–74. IEEE, March 2024.
- [3] Phang, J., M. Lampe, L. Ahmad, S. Agarwal, C. M. Fang, A. R. Liu, et al. "Investigating Affective Use and Emotional Well-Being on ChatGPT." *arXiv preprint arXiv:2504.03888*, 2025.
- [4] Khan, Z. A., B. Swaminathan, S. Prakash, S. Jayanthi, and R. Kumari. "ChatGPT and Sentiment Analysis: Unveiling Emotional Insights in Textual Data." *Procedia Computer Science* 259 (2025): 52–60.
- [5] Esh, M. "Sentiment Analysis in ChatGPT Interactions: Unraveling Emotional Dynamics, Model Evaluation, and User Engagement Insights." *Technical Services Quarterly* 41, no. 2 (2024): 160–74.
- [6] Kim, M. "Unveiling the E-Servicescape of ChatGPT: Exploring User Psychology and Engagement in AI-Powered Chatbot Experiences." *Behavioral Sciences* 14, no. 7 (2024): 558.
- [7] Rawat, P., M. Bajaj, S. Vats, and V. Sharma. "Redefining Human-AI Interactions: Unveiling ChatGPT's Profound Emotional Understanding." In *A Practitioner's Approach to Problem-Solving Using AI*, 1–19. Bentham Science Publishers, 2024.
- [8] Mishra, D., R. K. Mishra, and R. Agarwal. "Human–Computer Interaction and User Experience in the Artificial Intelligence Era." *Trends in Computer Science and Information Technology Research* (2025): 31–54.
- [9] Flores, H. I. O. "Creating an Adaptive Voice and Language Model Capable of Emotional Response and Self-Profiling to Emulate User Personality." *Ciencia Latina: Revista Multidisciplinar* 9, no. 1 (2025): 3454–71.
- [10] Grahm, C., and R. Stough. "Consumer Perceptions of AI Chatbots on Twitter (X) and Reddit: An Analysis of Social Media Sentiment and Interactive Marketing Strategies." *Journal of Research in Interactive Marketing* (2025).
- [11] Wu, Tianyu, Shizhu He, Jingping Liu, Siqi Sun, Kang Liu, Qing-Long Han, and Yang Tang. "A Brief Overview of ChatGPT: The History, Status Quo and Potential Future Development." *IEEE/CAA Journal of Automatica Sinica* 10, no. 5 (2023): 1122–36. <https://doi.org/10.1109/jas.2023.123618>.
- [12] Maroto-Gómez, Marcos, Sara Marqués-Villaroya, José Carlos Castillo, Álvaro Castro-González, and María Malfaz. "Active Learning Based on Computer Vision and Human–Robot Interaction for the User Profiling and Behavior Personalization of an Autonomous Social Robot." *Engineering Applications of Artificial Intelligence* 117 (January 2023): 105631. <https://doi.org/10.1016/j.engappai.2022.105631>.
- [13] Zhang, Hengjie, Xiaomin Wang, Weijun Xu, and Yucheng Dong. "From Numerical to Heterogeneous Linguistic Best–Worst Method: Impacts of Personalized Individual Semantics on Consistency and Consensus." *Engineering Applications of Artificial Intelligence* 117 (January 2023): 105495. <https://doi.org/10.1016/j.engappai.2022.105495>.
- [14] Li, Jing, Xinwei Zhang, Keqin Wang, Chen Zheng, Shurong Tong, and Benoit Eynard. "A Personalized Requirement Identifying Model for Design Improvement Based on User Profiling." *AI EDAM* 34, no. 1 (2020): 55–67. <https://doi.org/10.1017/S0890060419000301>.
- [15] Liu, Huan. "Feature Engineering for Machine Learning and Data Analytics." Edited by Guozhu Dong. March 2018. <https://doi.org/10.1201/9781315181080>.
- [16] Lee, Won-Jo. "A Study on Word Cloud Techniques for Analysis of Unstructured Text Data." *The Journal of the Convergence on Culture Technology* 6, no. 4 (2020): 715–20. <https://doi.org/10.17703/jcct.2020.6.4.715>.
- [17] Quinn, Kristen M., Xiaodong Chen, Louis T. Runge, Heidi Pieper, David Renton, Michael Meara, Courtney Collins, Claire Griffiths, and Syed Husain. "The Robot Doesn't Lie: Real-Life Validation of Robotic Performance Metrics." *Surgical Endoscopy* (October 2022). <https://doi.org/10.1007/s00464-022-09707-8>.
- [18] White, Jules, Quchen Fu, Sam Hays, Michael Sandborn, Carlos Olea, Henry Gilbert, Ashraf Elnashar, Jesse Spencer-Smith, and Douglas C. Schmidt. "A Prompt Pattern Catalog to Enhance Prompt Engineering with ChatGPT." *arXiv*, February 2023. <https://doi.org/10.48550/arxiv.2302.11382>.
- [19] Feng, Yunhe, Pradhyumna Poralla, Swagatika Dash, Kaicheng Li, Vrushabh Desai, and Meikang Qiu. "The Impact of ChatGPT on Streaming Media: A Crowdsourced and Data-Driven Analysis Using Twitter and Reddit." May 2023. <https://doi.org/10.1109/bigdatasecurity-hpsc-ids58521.2023.00046>.

- [20] Dewi, Christine, Rung-Ching Chen, Henocho Juli Christanto, and Francesco Cauteruccio. "Multinomial Naïve Bayes Classifier for Sentiment Analysis of Internet Movie Database." *Vietnam Journal of Computer Science* 10, no. 4 (2023): 485–98. <https://doi.org/10.1142/s2196888823500100>.
- [21] He, Biao, Danial Jahed Armaghani, and Sai Hin Lai. "Assessment of Tunnel Blasting-Induced Overbreak: A Novel Metaheuristic-Based Random Forest Approach." *Tunnelling and Underground Space Technology* 133 (March 2023): 104979. <https://doi.org/10.1016/j.tust.2022.104979>.
- [22] Aziz, Muhammad Maulidan, Mahendra Dwifabri Purbalaksono, and Adiwijaya Adiwijaya. "Method Comparison of Naïve Bayes, Logistic Regression, and SVM for Analyzing Movie Reviews." *Building of Informatics, Technology and Science (BITS)* 4, no. 4 (2023): 1714–20. <https://doi.org/10.47065/bits.v4i4.2644>.
- [23] Samih, Amina, Abderrahim Ghadi, and Abdelhadi Fennan. "Enhanced Sentiment Analysis Based on Improved Word Embeddings and XGBoost." *International Journal of Power Electronics and Drive Systems* 13, no. 2 (2023): 1827–36. <https://doi.org/10.11591/ijece.v13i2.pp1827-1836>.
- [24] Ikotun, Abiodun M., Absalom E. Ezugwu, Laith Abualigah, Belal Abuhaija, and Jia Heming. "K-Means Clustering Algorithms: A Comprehensive Review, Variants Analysis, and Advances in the Era of Big Data." *Information Sciences* 622 (2022). <https://doi.org/10.1016/j.ins.2022.11.139>.
- [25] Goloboff, Pablo A., and Martín E. Morales. "TNT Version 1.6, with a Graphical Interface for MacOS and Linux, Including New Routines in Parallel." *Cladistics* 39, no. 2 (2023). <https://doi.org/10.1111/cla.12524>.
- [26] Park, Sangil, and Jun-Ho Huh. "A Study on Big Data Collecting and Utilizing Smart Factory-Based Grid Networking Big Data Using Apache Kafka." *IEEE Access* 11 (2023): 96131–42. <https://doi.org/10.1109/ACCESS.2023.3305586>.
- [27] Ji, Hangxu, Su Jiang, Yuhai Zhao, Gang Wu, Guoren Wang, and George Y. Yuan. "BS-Join: A Novel and Efficient Mixed Batch-Stream Join Method for Spatiotemporal Data Management in Flink." *Future Generation Computer Systems* 141 (April 2023): 67–80. <https://doi.org/10.1016/j.future.2022.11.016>.
- [28] Nilashi, Mehrbakhsh, Rabab Ali Abumalloh, Masoumeh Zibarzani, Sarminah Samad, Waleed Abdu Zogaan, Muhammed Yousoof Ismail, Saidatulakmal Mohd, and Noor Adelyna Mohammed Akib. "What Factors Influence Students' Satisfaction in Massive Open Online Courses? Findings from User-Generated Content Using Educational Data Mining." *Education and Information Technologies* (March 2022). <https://doi.org/10.1007/s10639-022-10997-7>.
- [29] Li, Ying, R. K. Shyamasundar, and Xinheng Wang. "Special Issue on Computational Intelligence for Social Media Data Mining and Knowledge Discovery." *Computational Intelligence* 37, no. 2 (2021): 658–59. <https://doi.org/10.1111/coin.12457>.
- [30] Sun, Wenqi, Yuanjun Zhao, and Lu Sun. "Big Data Analytics for Venture Capital Application: Towards Innovation Performance Improvement." *International Journal of Information Management* (December 2018). <https://doi.org/10.1016/j.ijinfomgt.2018.11.017>.
- [31] Al-Shamri, Mohammad Yahya H. "User Profiling Approaches for Demographic Recommender Systems." *Knowledge-Based Systems* 100 (May 2016): 175–87. <https://doi.org/10.1016/j.knosys.2016.03.006>.
- [32] Hijikata, Yoshinori, Yuki Kai, and Shogo Nishida. "A Study of User Intervention and User Satisfaction in Recommender Systems." *Journal of Information Processing* 22, no. 4 (2014): 669–78. <https://doi.org/10.2197/ipsjip.22.669>.
- [33] He, Yi, Shixiong Cao, Yang Shi, Qing Chen, Ke Xu, and Nan Cao. "Leveraging Large Models for Crafting Narrative Visualization: A Survey." *arXiv* (Cornell University), January 2024. <https://doi.org/10.48550/arxiv.2401.14010>.
- [34] Ghemawat, Sanjay, Howard Gobioff, and Shun-Tak Leung. "The Google File System." In *Proceedings of the Nineteenth ACM Symposium on Operating Systems Principles – SOSP '03*. 2003. <https://doi.org/10.1145/945445.945450>.
- [35] Yin, Linzi, Jing Li, Zhaohui Jiang, Jiafeng Ding, and Xuemei Xu. "An Efficient Attribute Reduction Algorithm Using MapReduce." *Journal of Information Science* 47, no. 1 (2019): 101–17. <https://doi.org/10.1177/0165551519874617>.
- [36] Berros, Nisrine, Fatna El Mendili, Youness Filaly, and Younes El Bouzekri El Idrissi. "Enhancing Digital Health Services with Big Data Analytics." *Big Data and Cognitive Computing* 7, no. 2 (2023): 64. <https://doi.org/10.3390/bdcc7020064>.
- [37] Zhang, Zhan, Wenhao Li, Xiao Qing, Xian Liu, and Hongwei Liu. "Research on Optimal Checkpointing-Interval for Flink Stream Processing Applications." *Mobile Networks and Applications* 26, no. 5 (2021): 1950–59. <https://doi.org/10.1007/s11036-020-01729-7>.