

Study on Influencing Factors of Healing Music Records based on LightGBM

Jiayu Chen *

School of Informatics and Economics, Ningbo University of Finance and Economics, Ningbo, China

* Corresponding author email: 1922240003@s.nbufe.edu.cn

Abstract. At present, all kinds of music albums have become an indispensable product of spiritual consumption in public life. It can not only ease the pressure of life and work brought by external life, but also add some different interests to people's life. However, high-quality music is more popular with people, such as Dolby sound quality, and the demand for playing different types of high-quality music in different occasions and environments has greatly increased. This study will identify whether there is audience or artificial noise data during the recording of music albums, and explore the characteristics of various music albums, so that later music producers can better improve their music quality. To more specific, this study firstly collected the related dataset from Tianchi Platform. Then, data visualization e.g., correlation analysis is carried out to observe the data before passing it into the model. Subsequently, a typical machine learning algorithm called Light Gradient Boosting Machine is employed to train the dataset. The experimental results demonstrated the effectiveness of the model used in this study.

Keywords: Healing Records; Machine Learning; LightGBM; XGBoost.

1. Introduction

Music has many effects on people's emotions. For example, rock music can excite people's emotions, while healing music can relieve the pressure of current life and bring great healing to their emotions [1]. As a result, all kinds of songs are being released in large numbers, but healing songs will have a huge impact on people's mood modification, even change. This study will carry out exploratory data analysis on the records of the Healing Department, so as to analyze the influence of the recording process before the release of the record and its characteristics on the acceptance of the record itself and the audio and other characteristics. At the same time, this study will carry out machine learning to achieve the goal of classification [2].

In the early field of healing music, people mainly explored the role of music activity and music therapy in health care drawing, and examined the extent and type of music provision [3]. Besides, a paper traces the evolution of music as an alternative therapeutic option. However, most of current researches focused too much on medical treatment but ignored others e.g., mood and sleep status. They all extrapolated the function of healing music from the results, but did not have any scientific understanding of the music itself and specific characteristics. So, this study pays more attention to music research, including voice detection, intensity, loudness, audience popularity, etc. And through the study of these characteristics, healing music can be classified better in more detail.

This research focus on the Data visualization by Explorative Data Analysis (EDA) and Machine Learning- LightGBM to classify. This study would to show more comprehensive analysis on the popularity of the album and people in different mood will listen to different music. Finally, the Pearson coefficient is used to calculate the correlation coefficient between non character variables. When applying machine learning to classification, whether it has been considered to classify directly according to the prevalence, but the result is not ideal because of the insufficient data. To solve this problem, this study turned the focus on the explanation of the feature "liveness": if the liveness is greater than 0.8, it means that there is an audience when recording the album, and if the liveness is less than 0.8, there is no audience. Then 0.8 can be used as the dividing point; turn it into a binary problem, set the liveness as a tag, and set the value greater than 0.8 to 1, and the part less than 0.8 to 0.

At this point, this paper built a classification model to determine whether someone is recording an album through the influence factors of various albums. In order to improve the training loss convergence speed of the model and the accuracy of the model, this work conducted the StandardScaler mean and variance normalization for the training set and test set. Based on this, this study built the LGBMClassifier. Later, this study used optune to build the learning object, and set its learning direction to maximize as the guidance for model learning. The final accuracy of this model reaches 96%.

2. Method

2.1 Dataset Description and Preprocessing

This study used a discography dataset provided by TIANCHI platform, a dataset sharing platform of Aliyun [4]. The dataset is an analysis of The Cure album and provided us with healing records detailed information e.g., danceability, loudness, speechiness, liveness. This data can be used for exploratory data analysis and data visualization. The original dataset consists of 223 healing records' data and Including 23 features. At the beginning, this study discards the features of the column called 'Unnamed: 0'. At the same time, according to the prior knowledge analyzation, 'album_uri', 'album_img', 'track_uri', 'album_release_year' can be also discarded since these characteristics cannot be used as impact factors for later analysis. After preprocessing, there were 223 rows x 18 columns left in the dataset. The details of the dataset are shown in the Table 1.

Table 1. Some details of dataset

Feature	Item-1	Item-2	Item-3	Item-4
danceability	0.436	0.516	0.420	0.772
key_mode	G major	C major	C major	D major
loudness	-5.998	-5.872	-5.860	-9.788
liveness	0.1080	0.1360	0.0795	0.2820
track_popularity	33	28	35	34

2.2 Data Visualization

For track_name, for each album, their track_name is different. At the same time, each track_name reflects the characteristics of this album, so for the embodiment of track_name is very necessary and important. Word cloud will show the size of each keyword in the graph according to the number of occurrences of each keyword, so to the features of track_name, wordcloud is a good way to reflect the frequency of various tracks.

The study researched the relationship between album_name and album_popularity, because it is helpful for later visualization. Take these two features directly from the data df, and use the histogram to express the characteristic variables quantitatively.

According to various types of album_name's ratings of its corresponding danceability rhythm and energy music intensity and perception respectively account for the weight of all scores, and its weight is observed through the horizontal bar graph.

According to different tunes of all albums and their corresponding characteristic modes, calculate the category occupancy of various tunes. By setting the mode category as major or minor, respectively set the track_popularity to make and density estimation diagram to analyze various modes, and observe the relationship between tracks_popularity and its weight density.

Finally, this study used the Pearson coefficient to calculate the correlation coefficient between various non character variables, and then the part of the heatmap is used to express the in the Fig 1.

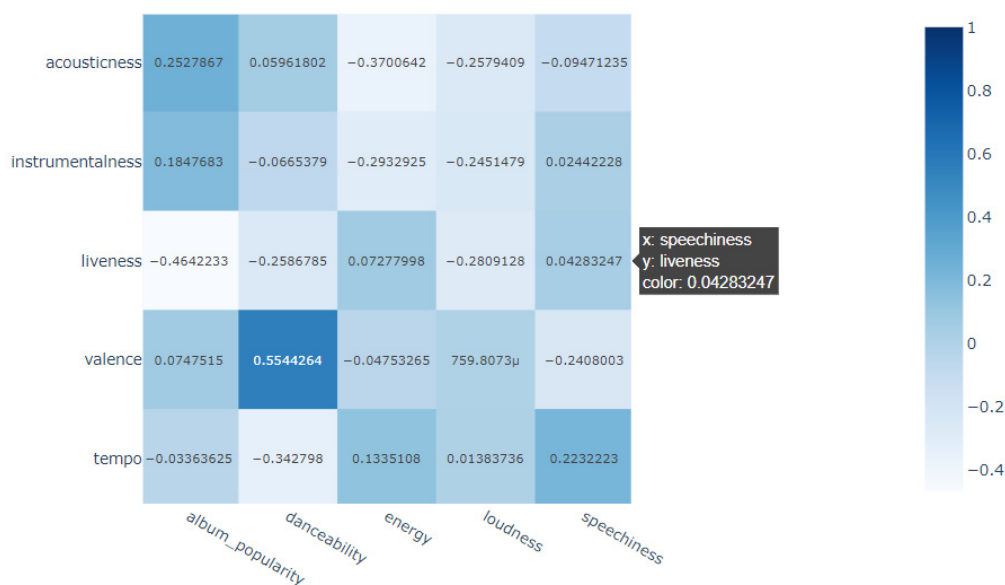


Fig 1. Part of the heatmap based on the features.

2.3 Machine Learning Model

The superiority of the machine learning models is proved in many prediction tasks before [5, 6]. Gradient Boosting Decision Tree (GBDT) is an enduring model in machine learning. Its main idea is to use weak classifiers (i.e., decision trees) to iteratively train to get the optimal model. This model has the advantages of good training effect and is not easy to over fit. GBDT is not only widely used in industry, but also used for multi classification, click through rate prediction, search sorting and other tasks. It is also a lethal weapon in various data mining competitions. According to statistics, more than half of the championship schemes in Kaggle competitions are based on GBDT [7].

Light Gradient Boosting Machine (LightGBM) is a framework for implementing GBDT algorithm [8-10]. It supports efficient parallel training, and has the advantages of faster training speed, lower memory consumption, better accuracy, supporting distributed and rapid processing of massive data. LightGBM adopts the Leaf wise growth strategy, which finds the leaf with the largest splitting gain from all the current leaves each time, then splits etc. Therefore, compared with Level wise, Leaf wise has the following advantages: under the same splitting times, Leaf-wise can reduce more errors and obtain better accuracy. The disadvantage of Leaf wise is that it may grow a deep decision tree and generate over fitting. Therefore, LightGBM will add a maximum depth limit to the Leaf wise to ensure high efficiency and prevent over fitting. The process is shown in Fig 2.

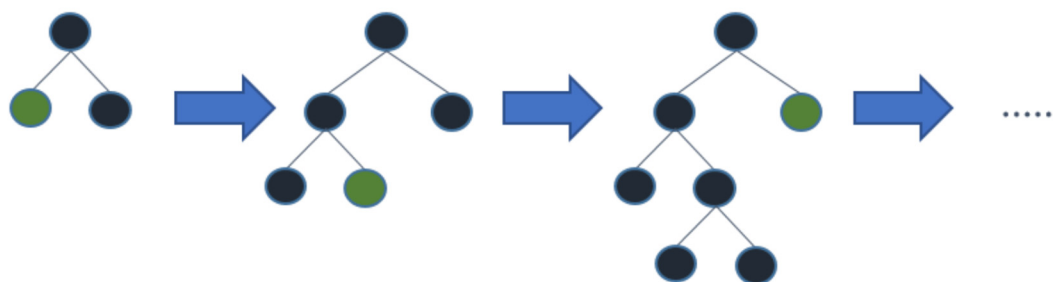


Fig 2. Leaf-wise tree growth

3. Result and Discussion

This study temporarily set the trend of the project as regression, take liveness as the analysis and research limitation of the dataset, and conduct a pairwise study on the influence relationship between the remaining features by drawing pair plot. It can be found from the above Fig 3 that when liveness=0 or liveness=1, it is closely related to other features. Therefore, if the threshold is set in liveness and the data is classified and transformed, the liveness can be transformed into a tag set for processing.

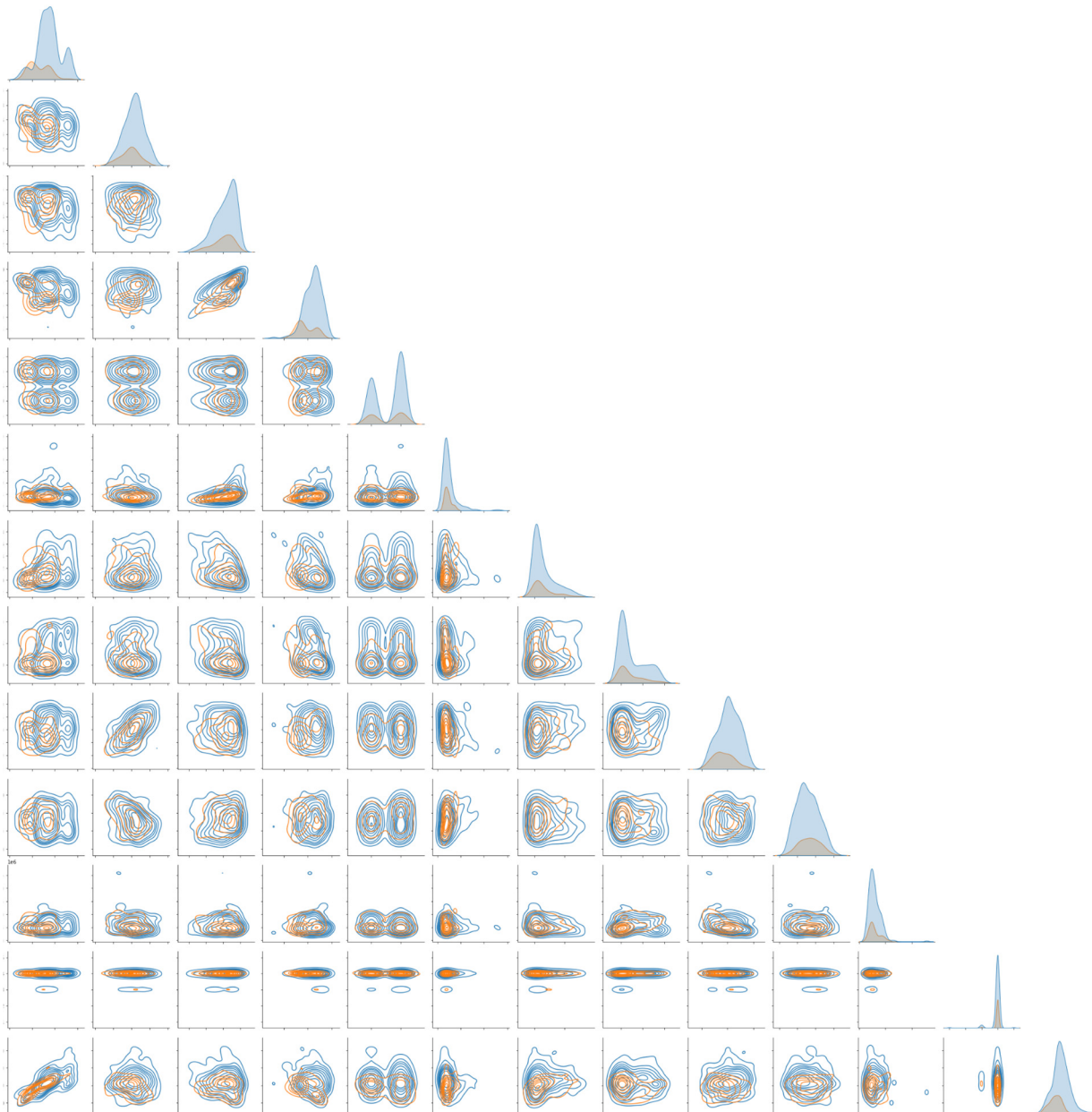


Fig 3. Pair plot of feature and influence

The data set can only observe the feature category and the impact of each feature on its category. If the data is transformed into supervised learning (or classification or regression problems) in machine learning, a set of tags is required as a supervisory signal to guide the model to train the feature set. Obviously, after intuitively observing the data set of this project, the data set of this project is missing which is defective for supervised learning. However, thinking backwards from the above results, if the relationship between liveness and other features is very close, and the interpretation of liveness in the dataset feature interpretation is "to detect whether there is a viewer in the record, and the value greater than 0.8 in liveness is likely to show that the orbit is in an active state". This study

takes 0.8 in liveness as a threshold, and convert the values greater than 0.8 and less than 0.8 into Boolean types, or 0 and 1. In this way, the project can be trend into a classification problem, and through this classification operation, the data of the test set can be predicted after the data set is divided.

	precision	recall	f1-score	support
0	0.95	1.00	0.97	39
1	1.00	0.67	0.80	6
accuracy			0.96	45
macro avg	0.98	0.83	0.89	45
weighted avg	0.96	0.96	0.95	45

Fig 4. Result of classify and forecast

After the model is automatically searched by optuna super parameters, a set of optimal super parameters is obtained. Putting this set of optimal super parameter sets into the model to train the training set data, and obtain a more mature model suitable for this data set type. Using this model to classify and forecast the data of the test set, and the results are shown in the Fig 4 above. Among them, the accuracy rate of 0 and 1 predicted by liveness is 95% and 100% respectively, the recall rate reaches 67% and 100%, and the accuracy is 96%.

Generally, the prediction and evaluation index of the model will not reach 100%, but the accuracy and recall rate of the model will reach 100%. According to the corresponding publicity: $\text{Recall} = \text{TP} / (\text{TP} + \text{FN})$, $\text{Precision} = \text{TP} / (\text{TP} + \text{FP})$, the number of test sets is 45, and the corresponding numbers of P and are 39 and 6, which are relatively small. If the relative order of magnitude of the data is small under the condition of good robustness of the model, it is easy to have an indicator of 100%. Here the accuracy that is equals to 96% can be taken as the final evaluation index of the model.

4. Conclusion

This study has conducted exploratory data analysis and visualization on the corresponding features of each album. At the same time, this study transformed the data before and after the threshold of liveness, and uses it as a label. After the data set is divided, a Light Gradient Boosting Model (lightGBM) classifier is built to transform the project into 'whether there is audience in the music recording process through various music characteristics of the album'. Before putting the data set into model training, this study used Optuna to find an optimal combination of super parameters suitable for the data set of this project under the super parameter adjustment scheme based on Bayesian super parameter optimization (TPESampler). This group of optimal super parameters was introduced into the model training, and the accuracy of the model finally reached 96%.

However, the precision and recall results of the model are not ideal. The main reason can be attributed to the lack of data sets. In the future, this problem can be solved to a certain extent by collecting and recording a large number of data sets. Finally, the model is used to evaluate the characteristics of the data set of this project in terms of 'feature importance', combined with the importance evaluation of features and the early analysis of feature data to conduct a comprehensive evaluation. The survival of popular music such as Kiss me Kiss me and wish can be speculated, and some of the music features it contains can be analyzed, so as to improve the quality of this category of music.

References

- [1] Fryberger A, Besada J L and Kanga Z. @ Newmusic# Soundart: Contemporary Music in the Age of social media. *Contemporary Music Review*, 41(4), 2022, 329-336.
- [2] Choo S, Yim H M, and Limb S J. Physical Music Records in the Age of Digital Music: Exploring Competitive Strategies of The Major Production. *Korea Business Review*, 2019, 23.2.
- [3] Daykin N, Bunt L, Mcclean S. Music and healing in cancer care: A survey of supportive care providers. *Arts in Psychotherapy*, 2006, 33(5):402-413.
- [4] Tianchi, The Cure discography, 2021, <https://tianchi.aliyun.com/dataset/93672>.
- [5] Mair, Carolyn, et al. An investigation of machine learning based prediction systems. *Journal of systems and software* 53.1, 2000, 23-29.
- [6] Yu Q, et al. Pose-guided matching based on deep learning for assessing quality of action on rehabilitation training. *Biomedical Signal Processing and Control* 72, 2022, 103323.
- [7] Guo Y Y, Ma H H, and Tian S Y. Quantitative investment model based on LightGBM algorithm. *Academic Journal of Computing & Information Science* 5.0.7.0, 2022.
- [8] Fan, J., Ma, X., Wu, L., Zhang, F., Yu, X., & Zeng, W. Light Gradient Boosting Machine: An efficient soft computing model for estimating daily reference evapotranspiration with local and external meteorological data. *Agricultural water management*, 225, 2019, 105758.
- [9] Taha, A. A., & Malebary, S. J. An intelligent approach to credit card fraud detection using an optimized light gradient boosting machine. *IEEE Access*, 8, 2020, 25579-25587.
- [10] Li, F., Zhang, L., Chen, B., Gao, D., Cheng, Y., Zhang, X., ... & Peng, J. A light gradient boosting machine for remaining useful life estimation of aircraft engines. In 2018 21st International Conference on Intelligent Transportation Systems (ITSC) (pp. 3562-3567), 2018, IEEE.