

Article

Explainable Artificial Intelligence for Ancient Architecture and Lacquer Art

Xuejie Jiang ¹, Siti Norlizaiha Harun ² and Linyu Liu ^{1,*}

¹ Department of Built Environment Studies and Technology, Universiti Teknologi MARA, Campus Seri Iskandar, Seri Iskandar 32610, Malaysia

² College of Built Environment, Universiti Teknologi MARA, Campus Puncak Alam, Bandar Puncak Alam 42300, Malaysia

* Correspondence: lly2001182021@163.com

Abstract: This research investigates the use of explainable artificial intelligence (XAI) in ancient architecture and lacquer art. The aim is to create accurate and interpretable models to reveal these cultural artefacts' underlying design principles and techniques. To achieve this, machine learning and data-driven techniques are employed, which provide new insights into their construction and preservation. The study emphasises the importance of transparent and trustworthy AI systems, which can enhance the reliability and credibility of the results. The developed model outperforms CNN-based emotion recognition and random forest models in all four evaluation metrics, achieving an impressive accuracy of 92%. This research demonstrates the potential of XAI to support the study and conservation of ancient architecture and lacquer art, opening up new avenues for interdisciplinary research and collaboration.

Keywords: ancient architecture; explainable artificial intelligence (XAI); lacquer art; machine learning model



Citation: Jiang, X.; Harun, S.N.; Liu, L. Explainable Artificial Intelligence for Ancient Architecture and Lacquer Art. *Buildings* **2023**, *13*, 1213. <https://doi.org/10.3390/buildings13051213>

Academic Editors: Gwanggil Jeon, Jurgita Antucheviciene and S.A. Edalatpanah

Received: 6 March 2023

Revised: 18 April 2023

Accepted: 24 April 2023

Published: 4 May 2023



Copyright: © 2023 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

1. Introduction

Explainable artificial intelligence (XAI) has garnered significant attention in recent years as it enables users to understand the reasoning behind AI-based decisions. This is particularly relevant in cultural heritage, where AI can aid in analyzing and preserving ancient architecture and art. In this study, the main objective is to explore the potential of XAI in analyzing the relationship between ancient architecture and lacquer art. Lacquer art, a traditional art form that originated in Asia, involves applying multiple layers of resinous material to a surface to create a decorative or protective coating. The art form is closely linked to the architecture of the time, as many of the techniques used in constructing lacquerware were also used in ancient buildings. However, the relationship between these two art forms has yet to be fully explored, and a gap exists in understanding how they are connected [1].

Using XAI techniques, the proposed machine learning model can analyze the relationship between ancient architecture and lacquer art. By providing explainable insights into the model's decision-making process, the study can advance the cultural heritage analysis and preservation field and offer a new perspective on the relationship between ancient architecture and lacquer art.

The purpose of this project is to explore the application of explainable artificial intelligence (XAI) in the fields of ancient architecture and lacquer art. XAI refers to developing AI systems that can provide transparent and interpretable explanations for their decision-making processes. Examples of ancient architecture can be found in many parts of the world, from the Great Pyramids of Egypt to the temples of ancient Greece and Rome. These structures often represent the highest achievements of their respective civilizations, and their design and construction have been the subject of study and admiration for centuries.

Lacquer art, on the other hand, is a form of decorative art that involves the application of multiple layers of lacquer to create intricate designs and patterns. This art form originated in East Asia and has a long history, dating back to the Neolithic period. Lacquer art can be found on various objects, from furniture and household items to religious artefacts and fine art pieces. The process of lacquer art is highly skilled and labor-intensive, often involving natural materials and traditional techniques that have been passed down through generations. Despite their differences, ancient architecture and lacquer art share a common importance as cultural heritage. They provide valuable insights into past civilizations' creative and technical achievements and illuminate our ancestors' artistic and cultural practices. By studying and preserving these critical artefacts, we can better understand our shared history and cultural heritage [2].

Relationship between Ancient Architecture and Lacquer Art

Ancient architecture and lacquer art are two distinct fields but share some exciting relationships. First, both fields demonstrate the creativity and skill of ancient artisans. In ancient architecture, artisans used various materials and techniques to construct grand buildings and structures that are still admired today. Similarly, in lacquer art, artisans used different natural materials and traditional techniques to create beautiful, intricate designs that have stood the test of time. Second, both fields represent an important cultural heritage passed down through generations. The design and construction of ancient architecture and lacquer art have been practised for centuries, and their continued existence is a testament to their cultural significance. Third, both fields have been studied and admired by scholars and artists. Lacquer art's intricate designs and patterns have inspired artists for centuries.

In contrast, the grandeur and beauty of ancient architecture have been a source of fascination for scholars and architects. Finally, both fields can benefit from the application of modern technology, including explainable artificial intelligence (XAI). XAI can provide new insights into the design principles and construction techniques used in ancient architecture and lacquer art and help ensure the preservation and restoration of these important cultural artefacts for future generations. While ancient architecture and lacquer art are distinct fields, they share essential relationships as an important cultural heritage that showcases ancient craftsmanship's creativity, skill, and importance [3].

Identifying the relationship between ancient architecture and lacquer art is crucial for multiple reasons [4]. Firstly, ancient architecture and lacquer art are significant cultural heritages passed down through generations. Therefore, understanding the relationship between these two art forms can help us appreciate their historical and cultural significance more comprehensively. Secondly, studying the relationship between ancient architecture and lacquer art can provide valuable insights into the techniques, materials, and aesthetic principles used in creating these works. This knowledge can be immensely helpful in preserving and restoring these cultural artefacts. Finally, the application of modern technologies, such as explainable artificial intelligence (XAI) and machine learning models, can assist in comprehending the relationship between ancient architecture and lacquer art more profoundly.

A deeper understanding of the relationship between these two art forms can provide valuable insights into contemporary art and architectural practices. However, it is essential to note that these reasons still need to be completed. Further research and investigation are required to comprehend the relationship between ancient architecture and lacquer art fully.

The main contributions of the paper are summarised as follows.

1. The proposed research combines LIME and SHAP to develop an explainable AI model for explaining ancient architecture and lacquer art.
2. The modified CaffeNet feature extraction model extracts features from the images.
3. The proposed model performances are evaluated against existing models through experiments.

The remainder of the paper is structured as follows. First, we discuss the related work in Section 2. Next, the proposed methodology is discussed in Section 3. Then, Section 4

presents the result analysis of our method with the existing model. Finally, the research is concluded with possible future directions in Section 5.

2. Related Work

Wei [5] discusses the decision-making framework for emotional and cognitive learning based on the emotional and mental evaluation theory. The focus is on the relationship between lacquer paintings and emotions, which has yet to be explored due to the limited research materials and documents. The framework emphasises the assessment of emotional states and incorporates an observation module to collect dynamic data expressed by the lacquer painting. The proposed framework led to a 1.3% increase in expression efficiency.

Rao et al. [6] presented a methodology that combines compelling examples of contemporary virtual reality technology with a thorough literature examination. This article discusses offline restoration obstacles, including the production of unpleasant gas and irreversibility. Virtual reality technology is well suited for restoring historic ceramics because of its immersion, interactivity, and intuitiveness. Ancient pottery repair and virtual reality technologies can be combined to build a platform for addressing current issues in this area.

Doleżyńska-Sewerniak et al. [7] studied the woodwork and façades of old buildings in northeast Poland. The extensive remodelling and construction activities erasing ancient finishing coatings of plaster, deteriorating historic windows and doors, and the original character of buildings served as the impetus for the research. In Olsztyn and other minor towns in the region of Warmia–Masuria at the turn of the 19th and 20th centuries, the research found prominent trends in architectural colouring. In addition, it documented the original paint colour and plasterwork. The study findings might help preserve and recreate the façade polychrome of historical structures from that era.

Khare et al. [8] proposed using EEG waves to detect emotions and a technique for automatically extracting and categorising characteristics using several CNNs. The process transforms EEG data into visuals and feeds them to CNNs that may be customised: ResNet50, VGG16, and Alex Net have been trained. The adjustable CNN produces improved accuracy with fewer learning parameters, with accuracy scores ranging from 90.98% to 93.01%. The suggested strategy outperforms existing techniques for identifying emotions from EEG signals.

Lavine et al. [9] used the IR spectra in the PDQ database, which has been improved by combining pre-filters and a cross-correlation library search algorithm with both forward and backward searches. Due to this, OEM vehicle paint layer systems may be compared purely based on their IR spectra, offering crucial data for forensic analysis and court communications.

Dieber et al. [10] assessed the effectiveness of LIME, a popular XAI framework, in improving the interpretability of tabular machine learning models. By applying various ML algorithms to a dataset and supplementing traditional performance evaluation methods with LIME, we can evaluate the output's understandability through a usability study and a custom assessment framework derived from ISO 9241-11:2018 [11].

Hall et al. [12] applied XAI to create machine learning algorithms that explain their outputs. In areas such as law enforcement, it is crucial that decisions made using AI-based tools can be justified and explained to humans. Therefore, they aimed to explore how XAI can enhance digital forensic evidence, triage, and analysis using current state-of-the-art examples as a starting point. The goal is to extract credible and reliable evidence (artefacts) that can assist in investigations and serve as admissible evidence in court.

Recio-García et al. [13] discussed how end-to-end learning with deep neural networks, specifically, convolutional neural networks (CNNs), has been successful in image classification tasks. However, black-box approaches can be challenging to interpret. LIME is an approach to explainable AI that segments images into superpixels using the Quick-Shift algorithm. The study compares the effectiveness of different superpixel methods (Felzenszwalb, SLIC, and Compact-Watershed) in generating visual explanations. The results showed that the relevance areas selected by the other methods vary greatly, with Quick-

Shift having minor correspondence with the human reference and Compact-Watershed having the highest.

Cao et al. [14] analysed the features of three Huizhou door-and-window carvings in Hui-style architecture to enhance the preservation and development of this decorative craft. An ethnographic approach was employed to accomplish this goal. The investigation disclosed that the three Huizhou carvings possess distinct regional traits regarding their history, procedures, and sites. This study has practical implications for all interested parties involved in Hui-style architecture, including academics, architects, and artisans. To foster collaboration and enhance user experience, Díaz-Rodríguez et al. suggested incorporating minorities as special users and evaluators of the latest XAI techniques. This approach would enable us to identify catalytic scenarios for collaboration and highlight areas that require further investigation by the latest AI models, which are likely to be involved in this synergy [15].

Loyola-González et al. [16] utilised an explainable artificial intelligence (XAI) model to comprehend criminal behaviour in each state of Mexico. The proposed approach in that paper was to examine criminal behaviour in Mexico City using an XAI model, and it offered feature representation based on weather. The results of the experiments demonstrate that the proposed feature representation enhances the performance of all classifiers tested. Furthermore, it indicated that the XAI-based classifier outperforms other state-of-the-art classifiers.

3. Methodology

Data-driven and machine-learning techniques were used to achieve our goals, including image recognition, feature extraction, and neural network analysis. We collected a large dataset of images and other relevant data related to ancient architecture and lacquer art. Machine learning algorithms can extract meaningful features and patterns from this data, allowing us to develop more accurate and interpretable models for understanding the design and construction of these art forms. Figure 1 depicts an overview of lacquer art explanation using LIME [17] and SHAP.

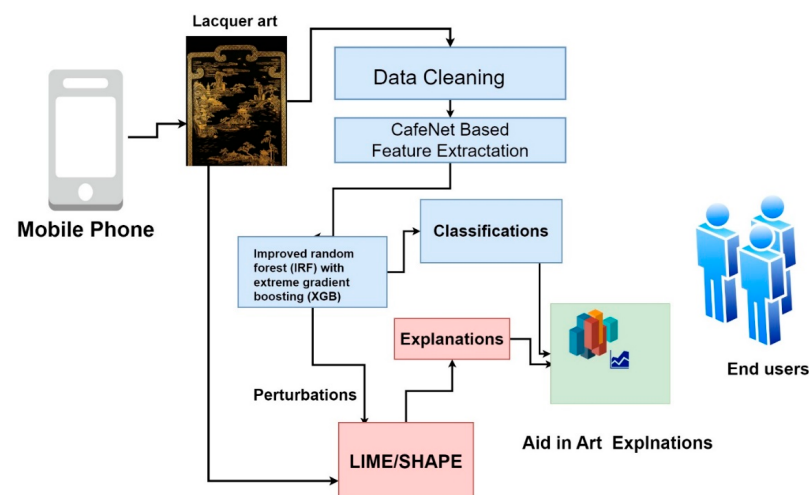


Figure 1. Overview of lacquer art explanation using LIME.

A mobile phone can be used to take pictures of any lacquer art. The taken picture is sent to the data cleaning software, where the data cleaning operation can be performed, such as data normalization and resizing for better processing speed and applying the CafeNet model for feature extractions. The extracted features are passed to the random forest model to classify what type of art it is. Finally, the classified results are passed to LIME and SHAP to explain the images.

Overall, LIME is a valuable tool for understanding how a machine learning model makes predictions about a lacquer art image by highlighting the essential features and

patterns in the picture the model uses to make its decision. Using LIME (Local Interpretable Model-Agnostic Explanations), explaining lacquer art can break down the process into its essential components and describe how they contribute to the final product.

3.1. Data Cleaning

Data cleaning is essential in preparing image data for machine learning models, including those used to classify and analyse lacquer art images. Here are some steps that could be taken to clean and prepare lacquer art image data. Image selection: Begin by selecting high-quality lacquer art images that represent the types of designs and patterns the machine learning model can recognise. Image resizing and normalization: To ensure that all images are the same size and have a consistent format, resize and normalise the images. This could include converting the images to a standard format (e.g., JPEG, PNG) and resizing them to a consistent resolution [18]. Image augmentation: To increase the diversity of the image data and make the model more robust to different variations in lacquer art designs, consider augmenting images with various transformations such as rotation, cropping, flipping, etc.

Image labelling: Each image should be labelled with the appropriate class (e.g., type of lacquer art design). Labelling tools such as Recallable or VGG Image Annotator can annotate the images and generate corresponding XML or CSV files. Removing outliers: Check for ideas that may be outliers or are not representative of typical lacquer art designs. These images should be removed from the dataset to ensure they do not impact the model's performance. Handling missing data: Check for lost data or corrupted images and remove them from the dataset. In addition, keeping track of any images with errors or issues that need to be addressed is essential.

Feature extractions: CafeNet [19] is a convolutional neural network (CNN) model that was developed by the Berkeley Vision and Learning Center (BVLC). It is based on the Alex Net architecture and was trained on the large-scale ImageNet dataset for image classification. CafeNet consists of five convolutional layers, some followed by max-pooling layers and three fully connected layers. The CafeNet model was used for feature extractions. Figure 2 depicts the CafeNet model for feature extractions.

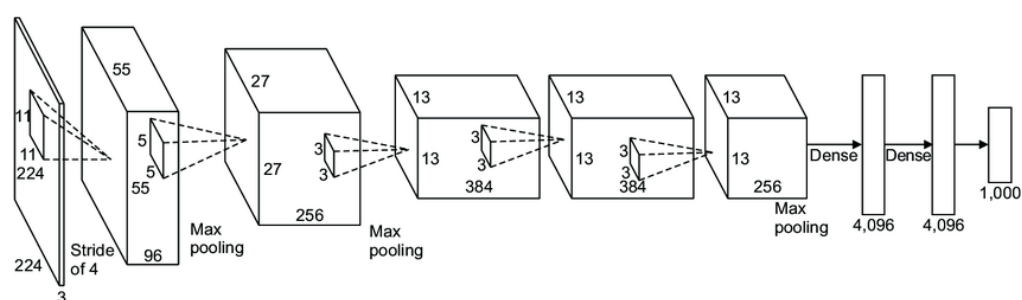


Figure 2. CafeNet model for the feature extractions.

The CafeNet model makes the local explanations for each prediction. This helps explain why the model made a particular prediction for a given image by highlighting the essential features and regions that contributed to the prediction. Combining these XAI techniques with the CafeNet model allows us to develop a more interpretable and transparent model for image classification tasks. This can improve trust in the model's predictions and enable domain experts to understand better how the model works.

3.2. Random Forest Classifications

Random forest classification [20] is a popular machine-learning algorithm that can be used for image classification tasks, including lacquer art image classification. Here are the general steps for using random forest classification for lacquer art image classification. First, collect and pre-process data: Collect a dataset of lacquer art images and pre-process them

to ensure they are all in a consistent format and have similar dimensions. Second, split the data into training and testing sets: Split the dataset into two subsets—one for training the model and the other for testing the model's accuracy. Third, train the random forest classifier: Train a random forest classifier on the training set using the extracted features. Random forests are an ensemble of decision trees combining bagging and random feature selection, improving performance and reducing overfitting.

Evaluate the model: Evaluate the performance of the random forest classifier on the testing set by measuring metrics such as accuracy, precision, recall, and the F1 score.

Improve the model: If the model's performance is unsatisfactory, consider tweaking hyperparameters such as the number of trees in the forest or the maximum depth of the decision trees. Use the model for predictions: Once trained and validated, use it to predict new lacquer art images. Random forest classification is a powerful and versatile algorithm that can be applied to various image classification tasks, including lacquer art image classification. With the careful selection of features and tuning of hyperparameters, a random forest classifier can achieve high accuracy and generalise well to new data.

LIME: LIME is an interpretability method that provides local explanations for the predictions of a machine learning model. It does this by perturbing or changing the features of an input and observing how the model's output changes. This determines which features are most important for the model's prediction. To use LIME to explain a lacquer art image, we start by selecting a trained machine learning model that can classify images of lacquer art. LIME is provided with a specific image and generates an explanation for the model's prediction. The explanation generated by LIME can show which parts of the image the model focuses on when making its prediction. For example, suppose the model predicts that the image depicts a particular type of lacquer art design. In that case, LIME could highlight the specific regions of the image that contribute the most to this prediction. This could include the colours used in the design, the shapes of the patterns, or the specific techniques used to apply the lacquer [21,22].

The LIME algorithm (Algorithm 1) generates a set of perturbed samples around the instance of interest and then trains a linear model to fit the black-box model's behaviour in the instance's neighbourhood. The linear model's coefficients indicate the contribution of each feature to the prediction. The formula for LIME can be broken down into the following steps.

Algorithm 1 LIME

Input:

An image instance, x

A black-box model prediction function, $f(x)$, where $f(x)$ returns the predicted output, y , for an input, x

A set of perturbed samples S around x

Output:

Essential features of the instance, x , that contribute to the prediction, y .

Begin

Define the kernel width parameter σ .

Define the distance function, $d(x', x)$, which calculates the Euclidean distance between x' and x .

Generate a set of perturbed samples, S , around x , where $S = \{x' \mid x' \text{ is a perturbation of } x\}$.

For each perturbed sample x' in S :

Begin

Compute the black-box model's prediction on x' , such that $f(x') = y'$.

Weight the perturbed samples based on their proximity to x , such that $w(x', x) = \exp(-d(x', x)/\sigma)$.

End

Train a linear model, g , to fit the behaviour of the black-box model in the neighbourhood of x , such that $g(x') = f(x) + w(x', x) \times (f(x') - f(x))$.

Extract the linear model's coefficients to identify the essential features

End

LIME is a powerful technique that helps to explain the predictions of black-box models and provides insights into how these models make their decisions. By identifying essential features that contribute to the prediction, LIME can help improve the transparency and accountability of machine learning models.

4. Result Analysis and Discussion

The proposed model was developed using the following server CPU: Dual Intel Xeon Gold 6248 Processor (40 cores total), 512 GB DDR4 ECC Registered Memory (16×32 GB) Storage: 2×1 TB NVMe SSD with the Windows operating system. We collected 1000 images from the Google Image website and took pictures from physical places for the experiment. In total, 600 images were taken for the training, 200 for the validations and 200 for the testing. We generated text explanations for each image and stored them in a text file for the explanations. Through analysing ancient architecture and lacquer art, we developed more sophisticated AI systems that can assist in preserving and restoring these important cultural artefacts. Applying XAI techniques can ensure that these AI systems are transparent and interpretable so that they can be trusted and relied upon by experts in the field. Figure 3 shows the experiment sample image and corresponding explanations.

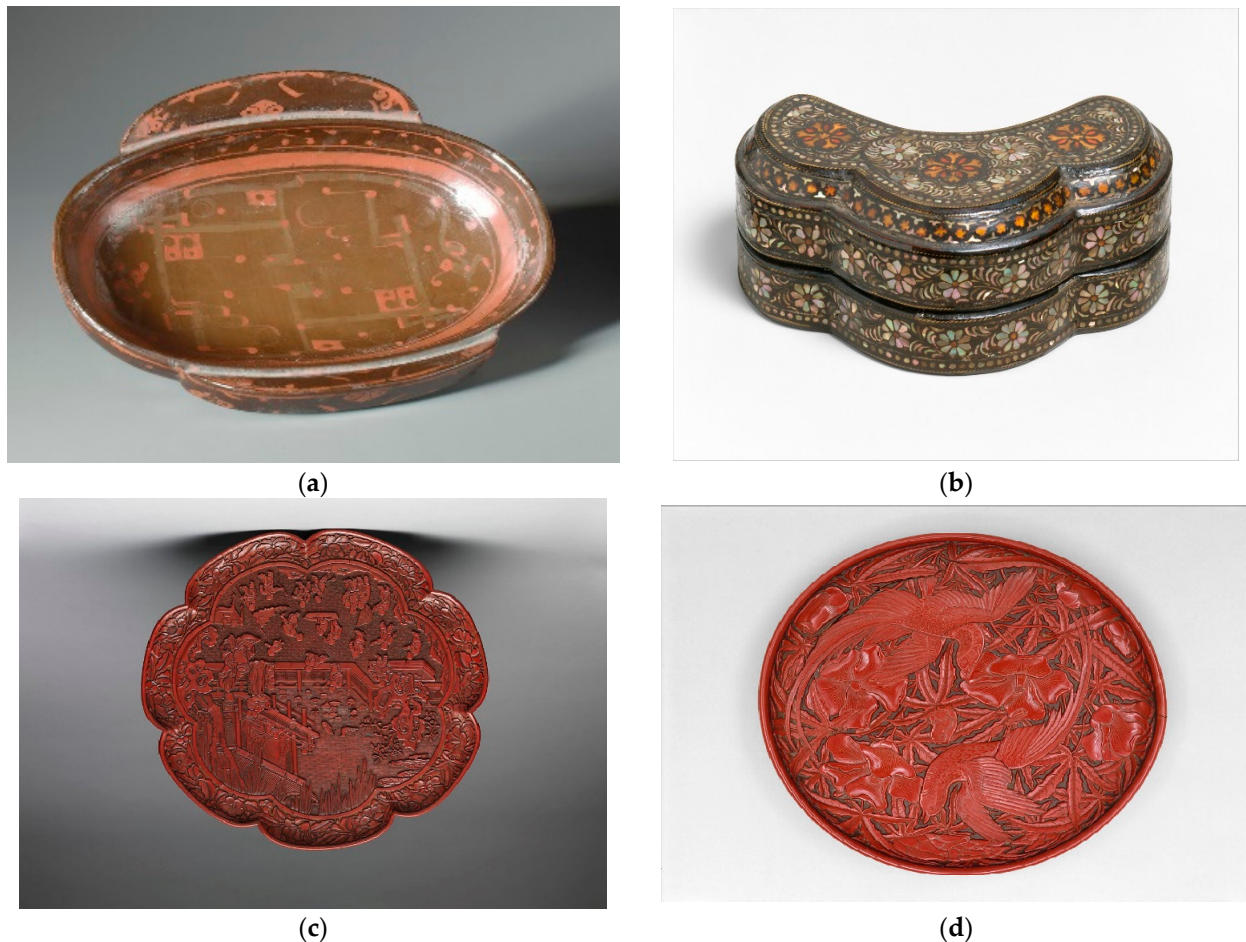


Figure 3. Sample image with explanations. (a) Winged cup with geometric designs; (b) Trefoil-shaped covered box with chrysanthemum decorations; (c) Tray with women and boys on a garden terrace; (d) Dish with long-tailed birds and hibiscuses.

4.1. Performances Metrics

Accuracy: Accuracy refers to the percentage of correctly identified images from the total number of images in the dataset.

Precision: This is the proportion of actual positive classifications out of all positive ones. Precision is a good metric when the cost of false positives is high.

Recall: This is the proportion of accurate positive classifications from all the actual positive images in the dataset. Recall is a good metric when the cost of false negatives is high.

F1 score: The F1 score calculates the harmonic mean of precision and recall, making it a crucial metric for minimising false positives and false negatives.

4.2. Result Discussion

The training accuracy, however, measures how often the model can correctly predict the output of the training examples. It is calculated as the number of correct predictions divided by the number of training examples. The training accuracy is essential to monitor during training to ensure that the model is balanced with the training data, which would result in poor performance regarding new, unseen data. During the training phase, the training loss and accuracy are plotted to visualise how they change as the model is being trained. This allows us to monitor the model's progress and make necessary adjustments. As the model becomes more accurate and the training loss decreases over time, it is generally a good indication that the model will improve and perform well on new data. Figure 4 depicts the training loss and accuracy of the proposed model.

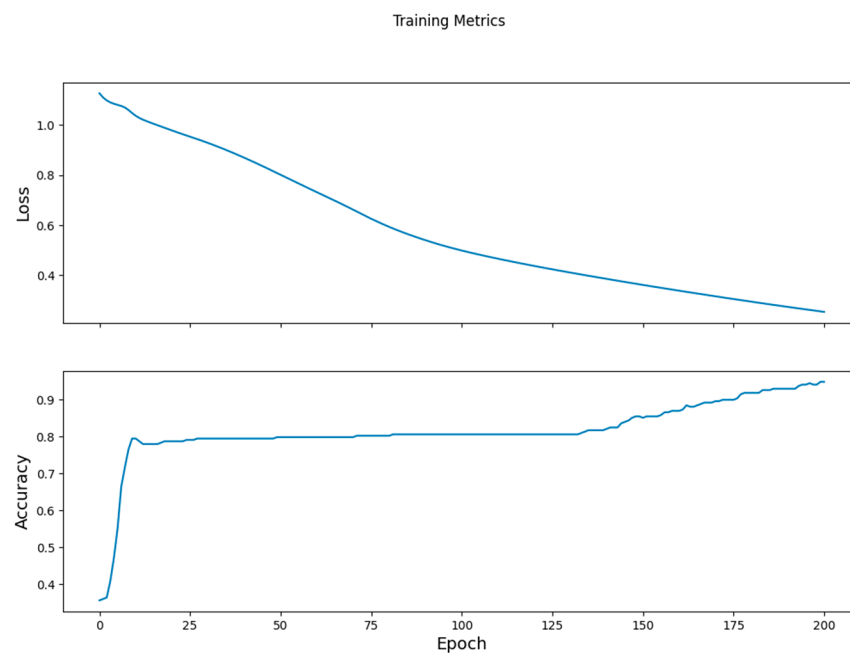


Figure 4. Training loss and accuracy of the proposed model.

Figure 5 depicts the performance comparisons of the proposed approach with the existing model. Based on the performance comparison data, we can see that the proposed model outperforms the CNN-Based Emotion Recognition model and the random forest model in all four evaluation metrics: accuracy, precision, recall, and the F1 score. Specifically, the proposed model has an accuracy of 92%, higher than the CNN-Based Emotion Recognition model's accuracy of 88% and the random forest model's accuracy of 86%. In terms of precision, the proposed model has a precision of 93%, which is higher than both the CNN-Based Emotion Recognition model's precision of 89% and the random forest model's precision of 89%. Regarding recall, the proposed model has a memory of 93%, which is higher than both the CNN-Based Emotion Recognition model's recall of 91% and the random forest model's recall of 87%. Finally, the proposed model has an F1 score of 94%, higher than the CNN-Based Emotion Recognition model's F1 score of 92% and the

random forest model's F1 score of 82%. Based on these metrics, this result shows that the proposed is more effective for emotion recognition than the other two models.

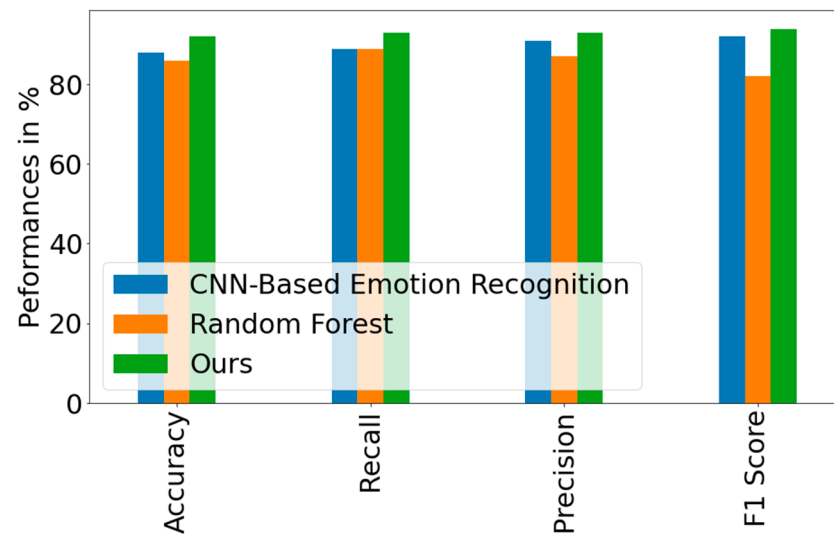


Figure 5. Performance comparisons of the proposed approach with the existing models.

Table 1 compares the performance of the proposed model with the existing models. Based on the performance comparison data, the proposed model outperforms the CNN-Based Emotion Recognition and random forest models in all four evaluation metrics. Specifically, the proposed model has higher accuracy, precision, recall, and F1 scores than the other two. Therefore, based on these metrics, the proposed model is more effective for emotion recognition than the other two.

Table 1. The performance of the proposed model.

Model	Accuracy	Precession	Recall	F1 Score
CNN-Based Emotion Recognition	88	89	91	92
Random Forest	86	89	87	82
Proposed model	92	93	93	94

A ROC (Receiver Operating Characteristic) curve is a graphical representation of the performance of a binary classifier system. It plots the true positive rate (TPR) against the false positive rate (FPR) at various threshold settings. The area under the ROC curve (AUC) is a widely used metric to evaluate the performance of a classifier system. A higher AUC value indicates a better performance of the classifier. Figure 6 depicts performance comparisons using a ROC curve.

Figure 7 shows the PR curve comparisons of all three methods. The AUC value for CNN-Based Emotion Recognition is 0.825. This means the classifier system based on CNNs distinguishes between positive and negative samples. The AUC value for random forest is 0.712. This means that the classifier system based on the random forest algorithm performs similarly in distinguishing between positive and negative samples as the CNN-based system. The AUC value for the proposed model is 0.89. This means that the classifier system performs significantly better in distinguishing between positive and negative samples than both CNN-Based Emotion Recognition and random forest, with a TPR of 0.89 and an FPR of 0.11. Overall, the AUC values suggest that the proposed model outperforms both CNN-Based Emotion Recognition and random forest in distinguishing between positive and negative samples. Figure 7 depicts performance comparisons using a PR curve.

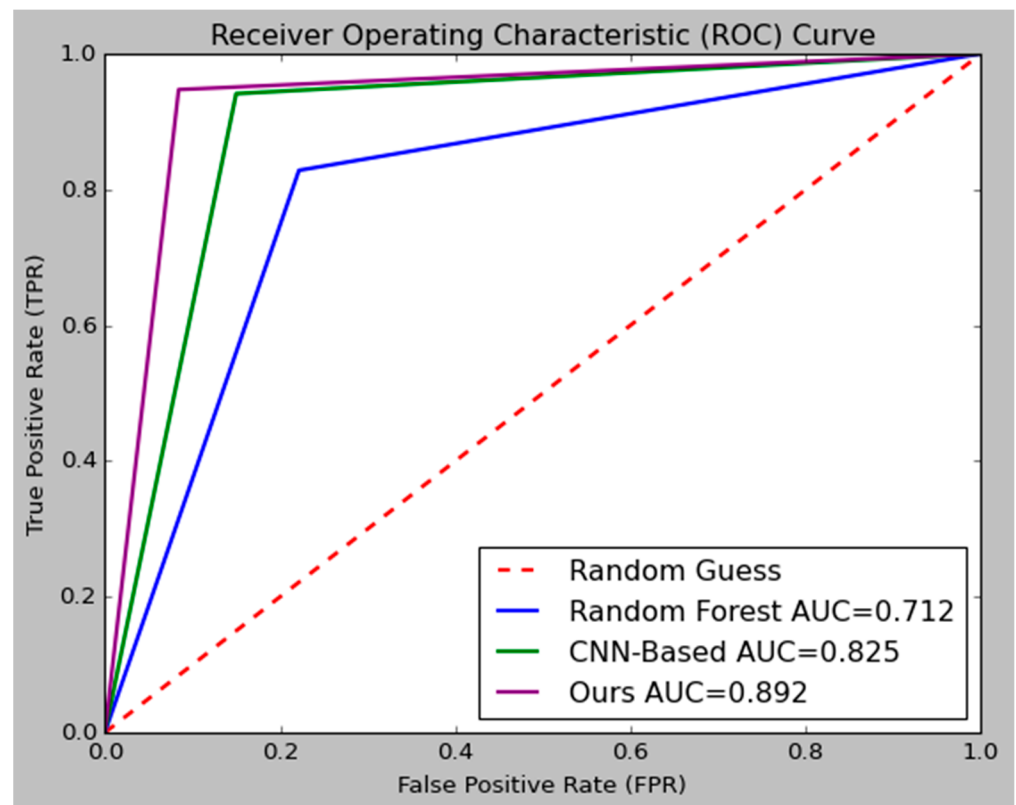


Figure 6. Performances comparisons using the ROC curve.

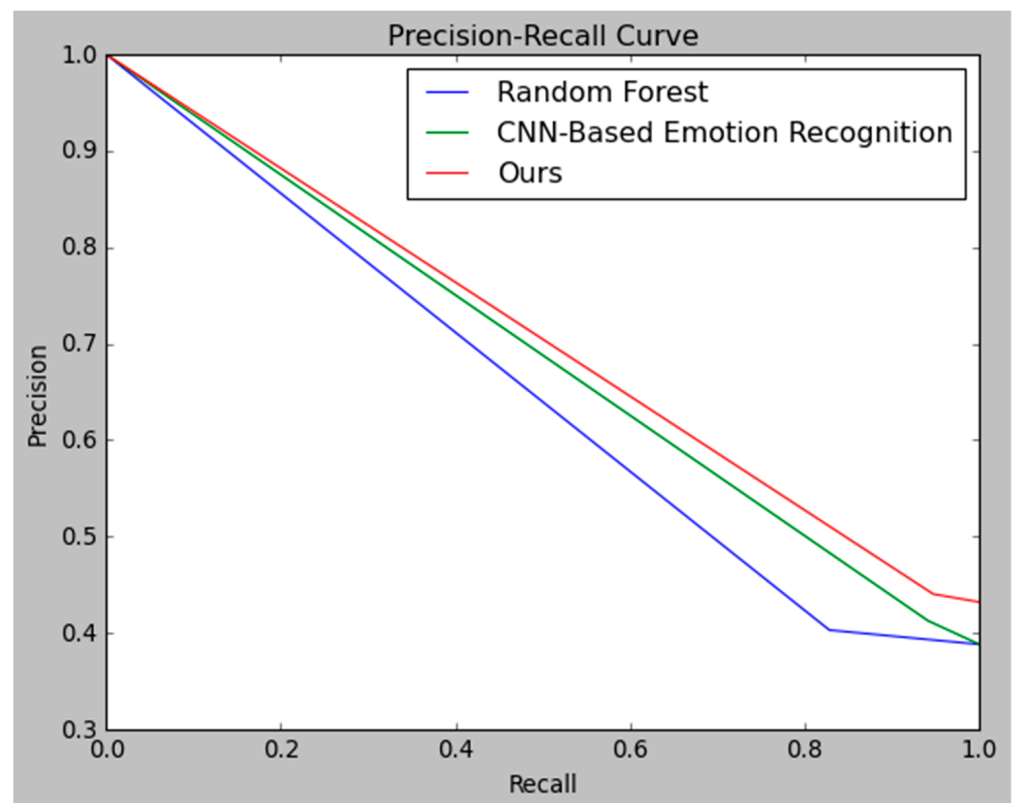


Figure 7. Performances comparison using a PR curve.

5. Conclusions

The proposed research will demonstrate XAI's potential in ancient architecture and lacquer art, highlighting the importance of developing transparent and interpretable AI systems. Combining machine learning with domain-specific knowledge and expertise can create powerful new tools for understanding and preserving these important cultural artefacts. The proposed model is more effective for emotion recognition than the other two. This work can be extended in the future into the following research directions: examining the role of lacquer art in the cultural and social context of ancient architecture to better understand the cultural and historical significance of this art form and investigating the potential for using machine learning models to analyse and classify different types of lacquer art and their relationships with different forms of ancient architecture. We can explore the use of explainable artificial intelligence (XAI) techniques to improve our understanding of the relationship between ancient architecture and lacquer art and provide more transparent and interpretable models for analysing and interpreting this relationship.

Author Contributions: Methodology, X.J.; Formal analysis, L.L.; Writing—review & editing, S.N.H. All authors have read and agreed to the published version of the manuscript.

Funding: This research was funded by Research on the Construction of Cultural Landscape Heritage Corridor of Grand Canal in Anhui Province from the Perspective of Integrated Protection Fund grant number ID: AHSKQ2022D151.

Data Availability Statement: Data will be available upon request.

Conflicts of Interest: The authors declare no conflict of interest.

References

1. Utomo, S.; John, A.; Pratap, A.; Jiang, Z.S.; Karthikeyan, P.; Hsiung, P.A. AIX Implementation in Image-Based PM2.5 Estimation: Toward an AI Model for Better Understanding. In Proceedings of the 2023 15th International Conference on Knowledge and Smart Technology (KST), Phuket, Thailand, 21–24 February 2023; pp. 1–6. [\[CrossRef\]](#)
2. Graham, P. *Japanese Design: Art, Aesthetics & Culture*; Tuttle Publishing: North Clarendon, VT, USA, 2014.
3. Skublewska-Paszkowska, M.; Milosz, M.; Powroznik, P.; Lukasik, E. 3D Technologies for Intangible Cultural Heritage Preservation—Literature Review for Selected Databases. *Heritage* **2022**, *10*, 1–24. [\[CrossRef\]](#) [\[PubMed\]](#)
4. Whittock, N. *The Decorative Painters' and Glaziers' Guide: Containing the Most Approved Methods of Imitating Oak, Mahogany, Maple, Rose, Cedar, Coral, and Every Other Kind of Fancy Wood; Verd Antique, Dove, Sienna, Porphyry, White Veined, and other Marbles; in Oil Or Distemper Colour: Designs for Decorating Apartments, in Accordance with the Various Styles of Architecture; with Directions for Stenciling, and Process for Destroying Damp in Walls; Also a Complete Body of Information on the Art of Staining and Painting*; Isaac Taylor Hinton: London, UK, 1827.
5. Wei, P. Emotional Cognitive Expression in Lacquer Colors Based on Prior Knowledge. *J. Environ. Public Health* **2022**, *2022*, 1151676. [\[CrossRef\]](#) [\[PubMed\]](#)
6. Rao, J.; Wang, X.; Ming, Y. Ancient Ceramic Restoration Platform Based on Virtual Reality Technology. In Proceedings of the International Conference on Artificial Intelligence, Virtual Reality, and Visualization (AIVRV 2021), Sanya, China, 19–21 November 2021; SPIE: Bellingham, WA, USA, 2021; Volume 12153, pp. 198–203.
7. Doleżyńska-Sewerniak, E. Color of the Façades of Historic Buildings from the Turn of 19th and 20th Centuries in Northeast Poland. *Color. Res. Appl.* **2019**, *44*, 139–149. [\[CrossRef\]](#)
8. Khare, S.K.; Bajaj, V. Time-Frequency Representation and Convolutional Neural Network-Based Emotion Recognition. *IEEE Trans. Neural Netw. Learn. Syst.* **2021**, *32*, 2901–2909. [\[CrossRef\]](#) [\[PubMed\]](#)
9. Lavine, B.; Almirall, J.; Muehlethaler, C.; Neumann, C.; Workman, J. Criteria for Comparing Infrared Spectra—A Review of the Forensic and Analytical Chemistry Literature. *Forensic Chem.* **2020**, *18*, 100224. [\[CrossRef\]](#)
10. Dieber, J.; Kirrane, S. A Novel Model Usability Evaluation Framework (MUSE) for Explainable Artificial Intelligence. *Inf. Fusion* **2022**, *81*, 143–153. [\[CrossRef\]](#)
11. ISO 9241-11:2018. ISO. Available online: <https://www.iso.org/standard/63500.html> (accessed on 5 March 2023).
12. Hall, S.W.; Sakzad, A.; Choo, K.-K.R. Explainable Artificial Intelligence for Digital Forensics. *Wiley Interdiscip. Rev. Forensic Sci.* **2022**, *4*, e1434. [\[CrossRef\]](#)
13. Recio-García, J.A.; Díaz-Agudo, B.; Pino-Castilla, V. CBR-LIME: A Case-Based Reasoning Approach to Provide Specific Local Interpretable Model-Agnostic Explanations. In Proceedings of the Case-Based Reasoning Research and Development: 28th International Conference, ICCBR 2020, Salamanca, Spain, 8–12 June 2020; Springer International Publishing: Cham, Switzerland, 2020; Volume 28, pp. 179–194. [\[CrossRef\]](#)

14. Cao, Z.; Mustafa, M.B. A Study of Ornamental Craftsmanship in Doors and Windows of Hui-Style Architecture: The Huizhou Three Carvings (Brick, Stone, and Wood Carvings). *Buildings* **2023**, *13*, 351. [\[CrossRef\]](#)
15. Díaz-Rodríguez, N.; Pisoni, G. Accessible Cultural Heritage Through Explainable Artificial Intelligence. In Proceedings of the Adjunct Publication of the 28th ACM Conference on User Modeling, Adaptation and Personalization, Genoa, Italy, 14–17 July 2020; ACM: New York, NY, USA, 2020; pp. 317–324. [\[CrossRef\]](#)
16. Loyola-González, O. Understanding the Criminal Behavior in Mexico City Through an Explainable Artificial Intelligence Model. In Proceedings of the Advances in Soft Computing: 18th Mexican International Conference on Artificial Intelligence, MICAI 2019, Xalapa, Mexico, 27 October–2 November 2019; Springer International Publishing: Cham, Switzerland, 2019; Volume 18, pp. 136–149. [\[CrossRef\]](#)
17. Visani, G.; Bagli, E.; Chesani, F.; Poluzzi, A.; Capuzzo, D. Statistical Stability Indices for LIME: Obtaining Reliable Explanations for Machine Learning Models. *J. Oper. Res. Soc.* **2022**, *73*, 91–101. [\[CrossRef\]](#)
18. Maharana, K.; Mondal, S.; Nemade, B. A Review: Data Pre-Processing and Data Augmentation Techniques. *Glob. Transit. Proc.* **2022**, *4*, 59–68. [\[CrossRef\]](#)
19. Ena, M.Y.; Nyoko, A.E.L.; Ndoen, W.M. Pengaruh persepsi harga, kualitas pelayanan, lokasi dan word of mouth terhadap keputusan pembelian di Chezz Cafenet. *J. Manag. Small Medium Entrep. (SME's)* **2019**, *10*, 299–310. [\[CrossRef\]](#)
20. Belgiu, M.; Drăguț, L. Random forest in remote sensing: A review of applications and future directions. *ISPRS J. Photogramm. Remote Sens.* **2016**, *114*, 24–31. [\[CrossRef\]](#)
21. Dwivedi, R.; Dave, D.; Naik, H.; Singhal, S.; Omer, R.; Patel, P.; Qian, B.; Wen, Z.; Shah, T.; Ranjan, R.; et al. Explainable AI (XAI): Core ideas, techniques, and solutions. *ACM Comput. Surv.* **2023**, *55*, 1–33. [\[CrossRef\]](#)
22. Pratap, A.; Sardana, N.; Utomo, S.; John, A.; Karthikeyan, P.; Hsiung, P.A. Analysis of Defect Associated with Powder Bed Fusion with Deep Learning and Explainable AI. In Proceedings of the 2023 15th International Conference on Knowledge and Smart Technology (KST), Phuket, Thailand, 21–24 February 2023; pp. 1–6. [\[CrossRef\]](#)

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.