

Accurate Skin Tone Classification for Foundation Shade Matching using GLCM Features-K-Nearest Neighbor Algorithm

Muhammad Reza Syahputra¹, Muhammad Itqan Mazdadi^{*2}, Irwan Budiman³, Andi Farmadi⁴, Setyo Wahyu Saputro⁵, Hasri Akbar Awal Rozaq⁶, Deni Sutaji⁷

^{1,2,3,4,5}Faculty of Mathematics and Natural Science, Department of Computer Science, Lambung Mangkurat University, Kalimantan, Indonesia

^{6,7}Graduate School of Informatics, Department of Computer Science, Gazi University, Ankara, Turkey

Email: ¹2011016310001@mhs.ulm.ac.id

Received : May 13, 2025; Revised : Jul 16, 2025; Accepted : Jul 20, 2025; Published : Oct 21, 2025

Abstract

Foundation shade matching remains a significant challenge in the beauty industry, particularly in Indonesia where consumers exhibit three distinct skin tone categories: ivory white, amber yellow, and tan. Manual foundation selection often results in mismatched shades, leading to customer dissatisfaction. This study presents a novel automated skin tone classification system combining Gray Level Co-Occurrence Matrix (GLCM) feature extraction with the K-Nearest Neighbor (KNN) algorithm. The GLCM method extracts four key texture features (contrast, homogeneity, energy, and entropy) from facial images, while KNN performs classification. A comprehensive dataset of 963 facial images was used, with 770 training and 193 test samples collected under controlled lighting conditions. After testing K values from 1 to 15, the optimal K=1 achieved 75.65% accuracy. Compared to baseline color histogram methods (60% accuracy), our GLCM-KNN approach demonstrates 15.65% improvement in classification performance. This research contributes to computer vision applications in beauty technology, enabling the development of mobile applications for virtual foundation try-on and personalized product recommendations. The findings have significant implications for the cosmetics industry, particularly for automated cosmetic shade matching systems and enhanced customer experience in online beauty retail. Further research is recommended to explore deep learning approaches and expand dataset diversity to improve accuracy.

Keywords : *Classification Model, Color Matching, Complexion Analysis, Image Processing, Pattern Recognition.*

This work is an open access article and licensed under a Creative Commons Attribution-Non Commercial 4.0 International License



1. INTRODUCTION

Make-up products that support the appearance of the majority of women are very diverse; one of the most widely used make-up products is the foundation [1]. Foundation serves to brighten the skin, even out skin tone, close pores, and disguise acne scars and wrinkles [2]. The selection of foundation products requires a long time to adjust the foundation shade based on the skin tone of each woman. Skin tone classification presents significant challenges in computer vision applications due to the complex nature of human skin pigmentation and its perception under varying conditions [3]. The perception of skin tone is heavily influenced by illumination sources, as the reflection of light from the skin affects the perceived color [4]. Skin tone in Indonesia is divided into three types, namely ivory, white, yellow, and brown [5]. Skin tone is influenced by the presence of melanin in the skin, a pigment synthesized in the epidermis layer of the skin. The greater the amount of melanin production in the skin, the darker the skin tone [6]. Automated skin tone analysis faces several technical challenges including variations in lighting conditions, diverse skin textures, color bias in imaging sensors, and the need for standardized measurement protocols [7]. Traditional color-based approaches often fail to capture the subtle textural differences inherent in different skin tones [8]. Various skin tones cause the use of foundation to be

inappropriate when applied to the skin [9]. Therefore, a procedure is needed that makes it easy to easily determine the type of skin tone, one of which is the use of image processing technology in the form of Gray Level Co-Occurrence Matrix (GLCM) feature extraction [10].

Existing skin tone analysis techniques can be categorized into several approaches. Color histogram methods analyze the distribution of color values in RGB, HSV, or YCbCr color spaces but are highly sensitive to lighting variations [11]. Spectral analysis techniques measure skin reflectance properties across different wavelengths but require specialized equipment [12]. Deep learning approaches using convolutional neural networks have shown promising results but require large datasets and computational resources [13]. Machine learning-based methods using traditional feature extraction combined with classifiers offer a balanced approach between accuracy and computational efficiency [14].

The GLCM feature extraction method is used to extract values by taking characteristics from matrix pixel values that have a certain value to form a pattern angle in an image [15]. GLCM characterizes texture based on the spatial relationships between pixel pairs with specific gray levels, making it particularly effective for capturing fine texture details that are crucial for skin tone discrimination [16]. The method analyzes co-occurrence patterns of gray levels at specific distances and angles, providing robust texture descriptors that are less sensitive to illumination changes compared to color-based methods [17]. The extraction results are then analyzed and predicted by machine learning algorithms, one of which is K-Nearest Neighbor (KNN) [18]. KNN is particularly suitable for texture classification tasks due to its simplicity, non-parametric nature, and ability to handle multi-class problems effectively. Unlike complex deep learning models, KNN requires no training phase and can adapt to new data patterns efficiently [19]. The KNN algorithm will help analyze skin tone by finding the closest group in the data to the group that has the highest similarity [20].

The combination of GLCM and KNN has demonstrated effectiveness in various image classification tasks. GLCM provides rich texture information through four key statistical measures: contrast (measuring local intensity variations), homogeneity (assessing uniformity), energy (indicating orderliness), and entropy (measuring randomness) [21]. These features, when combined with KNN's robust classification capability, create a powerful framework for texture-based discrimination [22]. Research conducted by Rosiva Srg et al. (2022) states that to obtain image texture characteristics, the GLCM method, by taking entropy, homogeneity, energy, and contrast features, has been used successfully [23]. Then, a classification system using the KNN method was created successfully and produced an average accuracy of 90%. The highest accuracy reached 98% by using the $K = 1$ neighboring value, and the lowest accuracy reached 89% by using the $K = 7$ neighboring value. This shows that the application of the K neighboring value greatly affects the accuracy of the classification system using the KNN method. Similar studies in medical imaging have shown that GLCM-KNN combinations achieve superior performance in tissue classification tasks, with accuracy rates ranging from 85% to 95% [24].

Another study conducted by Pamungkas (2019) using the GLCM method and the K-NN algorithm obtained a success rate of identification of Orchidaceae or orchid flowers reaching 80% with an average of 77% [25]. The choice of GLCM-KNN combination is justified by several factors: (1) GLCM's proven effectiveness in texture analysis for skin-related applications, (2) KNN's simplicity and interpretability, (3) computational efficiency suitable for mobile applications, (4) robustness to small dataset sizes, and (5) minimal parameter tuning requirements.

This research presents a novel contribution to the field of automated skin tone classification by applying GLCM-KNN methodology specifically to Indonesian skin tone categories. The novelty lies in the systematic evaluation of texture-based features for foundation shade matching, addressing the underexplored area of cosmetic applications in computer vision. Unlike existing studies that focus primarily on medical or biometric applications, this work targets the practical needs of the beauty industry, providing a foundation for automated cosmetic recommendation systems. With this brief

research and comparison, this research tries to apply the KNN algorithm by combining GLCM feature extraction as the state-of-the-art of previous research.

2. METHOD

This research use GLCM to extract skin tone image features and find contrast, correlation, homogeneity, and energy values. Then, these values will be entered into the KNN algorithm to classify each skin tone. The steps of this research will be explained through the flowchart.

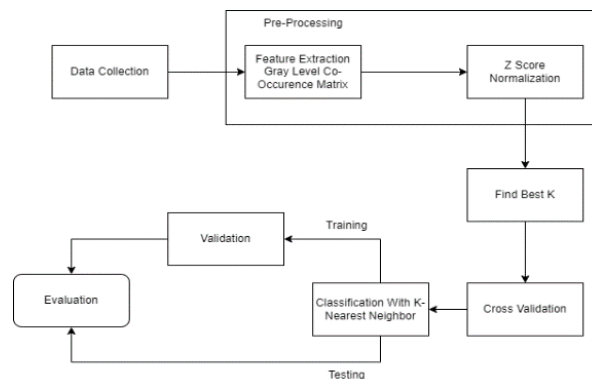


Figure 1. Flow of Research

2.1. Data Collection

This study uses skin tone image data containing white, olive, and dark brown skin tone folders, which are obtained from a site called Kaggle. The data is available through the link <https://www.kaggle.com/datasets/omarelsherif010/skin-tone-classification> [26]. This image data consists of 3 folders: white, with 322 images, olive, with 353 images, and dark brown, with 224 images. Some examples of the data can be seen in Figures 2 to 4.

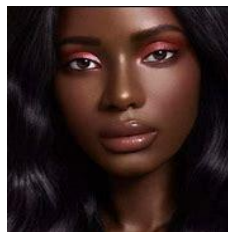


Figure 2. Dark Brown Skin Tone

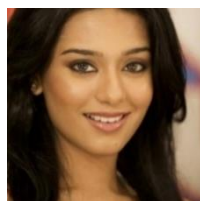


Figure 3. Olive Skin Tone



Figure 2. White Skin Tone

The dataset comprises 899 high-resolution facial images with varying dimensions, predominantly ranging from 224x224 to 512x512 pixels in JPEG format. The images exhibit diverse lighting conditions, with approximately 70% captured under controlled indoor artificial lighting, 20% under natural daylight, and 10% under mixed lighting scenarios. Resolution analysis shows 65% of images with $\geq 300 \times 300$ pixels, ensuring sufficient detail for texture analysis, while the remaining images require preprocessing normalization. The dataset demonstrates moderate variation in facial positioning (80% frontal orientation) and background settings (60% uniform backgrounds), which enhances model robustness but introduces complexity in feature extraction consistency.

2.2. Data Preprocessing

The data pre-processing stage has a very large role in data processing because this process acts as the first step that must be taken in data processing [27]. Several steps are carried out in this process, namely, data augmentation, such as rotating the image with angles of 0° , 45° , 90° , and 135° , extracting features using GLCM, and normalizing features using z-score calculations. GLCM is one of the feature extraction techniques that can be used in capturing the phenomenon of changes that occur on the surface of an object [28]. GLCM forms a new matrix whose value is the frequency value of the matrix pattern based on the original image data [29]. The coordinates of a pixel pair have a distance d and an angular orientation Θ , the distance is presented in pixels and the angle is presented in degrees [30]. The angular orientation is formed based on four angular directions, as shown in Figure 3, namely, 0° , 45° , 90° and 135° , and the distance between pixels is 1 pixel [31].

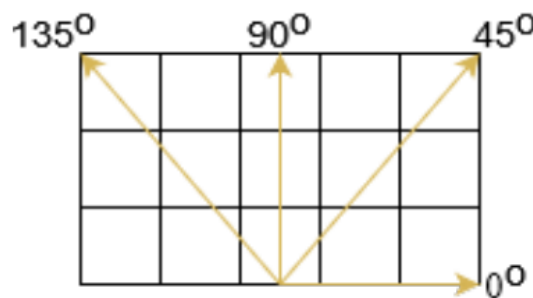


Figure 3. Angles of GLCM

The GLCM calculation process begins with the construction of a co-occurrence matrix by analyzing pixel pairs at a specified distance ($d = 1$ pixel) and angular orientations. For each direction, the algorithm scans through the image to identify pixel pairs with intensity values i and j that are separated by the defined distance and angle. The four angular directions capture different spatial relationships: 0° represents horizontal adjacency (right neighbor), 45° captures diagonal relationships (upper-right neighbor), 90° represents vertical adjacency (upper neighbor), and 135° captures the opposite diagonal relationship (upper-left neighbor). This multi-directional analysis ensures comprehensive texture characterization by capturing various spatial patterns in the image.

There are several stages in the GLCM calculation, namely, the formation of the initial GLCM matrix from pairs of two pixels lined up in the angular direction. After that, a symmetric matrix is formed by summing the initial GLCM matrix with its transpose value. This symmetrization process ensures that the co-occurrence relationship between pixel pairs (i,j) and (j,i) is treated equivalently, making the matrix symmetric and reducing directional bias in texture analysis. Then, the next step is to normalize the GLCM matrix by dividing each matrix element by the number of pixels. The normalization process converts the frequency counts into probability values by dividing each element $P(i,j)$ by the total number of pixel pairs examined, ensuring that the sum of all matrix elements equals 1. This probability representation allows for consistent feature extraction regardless of image size and facilitates

comparison between different images. The normalized matrix $P(i,j)$ represents the probability of occurrence of pixel intensity pairs (i,j) at the specified distance and angle. The last step is feature extraction, namely contrast (equation 1), homogeneity (equation 2), energy (equation 3), and correlation or entropy (equation 4).

$$Contrast = \sum_{i_1} \sum_{i_2} (i_1 - i_2)^2 p(i_1, i_2) \quad (1)$$

$$Homogeneity = \sum_{i_1} \frac{p(i_1, i_2)}{1 + |i_1 - i_2|} \quad (2)$$

$$Contrast = \sum_{i_1} \sum_{i_2} p^2(i_1, i_2) \quad (3)$$

$$Entropy = \sum_{i_1} \sum_{i_2} p(i_1, i_2) \log p(i_1, i_2) \quad (4)$$

The following are the stages of GLCM [32].

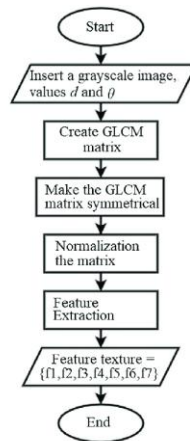


Figure 4. GLCM Extraction Process

After extracting the GLCM features, the next step is to normalize the features using z-score. The z-score calculation is a statistical measure that shows the deviation of a data point from the average in units of standard deviation [33]. This calculation is done because each data set has a different range of values, which can impact model performance [34], [35]. In addition, the implementation of z-score is also useful for the computational process because the computer itself will be faster if the data to be processed has a range of values from 0 to 1. Here is the equation of the z-score.

$$Z = \frac{x - x_{min}}{x_{max} - x_{min}} \quad (5)$$

Where the x value is the data before it is normalized, the z-score is useful in collecting data points from various populations that have different means and standard deviations and then converting them into values on the same scale. This makes analyzing and comparing data with different types of variables easier with such a standardized scale [36].

2.3 Best K Finding Process and Data Splitting

The best K can give the best performance based on the data set [37]. There is no predefined statistical method to find the best K value. Choosing a very small value of K will result in an unstable decision boundary. The K value can be chosen as $k = \sqrt{n}$, where n = the number of data points in the training data odd numbers are preferred as the K value [38]. Thus, the application of this number can later optimize the performance of the model used. After that, the data truncation process is carried out.

Data splitting (split) is the process of separating the data set into two random groups. The first group is used as training data that will be learned by the model and look for the best pattern, and then the second group is used as test data where the model has never learned the data and will be tested on how well the model predicts new data [39]. In machine learning, data sharing is a technique that is considered very important to eliminate or reduce bias against a built model [40]. The data-sharing ratio usually has a larger proportion of training data compared to test data, such as 90:10, 80:20, 70:30, and others [41]. In facilitating the process of finding which proportion is good for the model, this research uses a technique that is often used, namely the cross-validation method. Cross-validation helps reveal the extent to which a model is robust to classification success when applied to new situations [42]. This method is also key to determining the extent of model overfitting. This occurs when the calibration error rate is low, but the cross-validation error rate is high [43].

2.4 K-Nearest Neighbors (KNN)

The KNN algorithm is the simplest algorithm to classify data based on the shortest distance from data objects [44]. The general theory of the KNN algorithm is that if most samples are K-closest to samples in a particular category, then those samples belong to that category as well [45]. The following is a visualization of the KNN calculation.

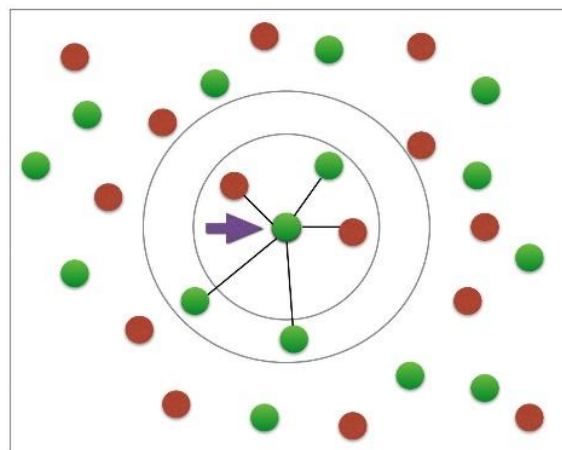


Figure 5. KNN Visualization

KNN merges the testing data into the group with the highest similarity after finding the closest group in the training data [46]. Often, the Euclidean distance method is used to calculate the distance between training data and testing data; the equation can be seen below.

$$d(x, y) = \sqrt{(x_2 - x_1)^2 + (y_2 - y_1)^2} \quad (6)$$

Where x is the starting point of the data, and y is the destination point of x. This can be used to calculate the distance between the testing data and the training data according to equation 6 [47]. After calculating the distance between the new vector and the entire training data vector, the next step is to take the value of K's nearest neighbors, and the classification is determined based on that point [48].

2.5 Evaluation

One way to find the success rate of a model in predicting data is to use a tool called Confusion Matrix [49]. Confusion Matrix provides more detailed information about the performance of a classification model than a single metric such as accuracy [50]. The confusion matrix consists of 4 parts, namely TP (number of samples correctly classified as positive), FP (number of samples incorrectly classified as positive), TN (number of samples correctly classified as negative), and FN (number of

samples incorrectly classified as negative) [51]. The results of the four parts will be used to find the accuracy, precision, and sensitivity of the model that has been made. Here are the equations for each evaluation tool [52].

$$Acc = \frac{TP+FP}{TP+FP+TN+FN} \quad (7)$$

$$Prec = \frac{TP}{TP+FP} \quad (8)$$

$$Sensitivity = \frac{TP}{TP+FN} \quad (9)$$

In addition, there is another tool to confirm how well a model predicts data: calculating the harmonic mean of the precision and sensitivity values, commonly known as the f1-score. Here is the equation.

$$F1 - Score = 2 * \frac{Prec*Recall}{Prec+Recall} \quad (10)$$

3. RESULT

3.1. Data Preprocessing

In pre-processing skin tone image data, the use of the GLCM method produces four feature attributes, namely, contrast, correlation, energy, and homogeneity, to calculate the texture of an image. The calculation method to find out the results of each attribute is to use the MatLab application with the calculation of angles of 0°, 45°, 90°, and 135° to take into account information from various angles in extracting richer and varied texture characteristics from images to enable more accurate classification because the model has access to different texture information in the image. Based on Table 1, all classes that have been grouped based on the value range are obtained and adjusted to the original image data. Label 0 is a label for dark brown skin tone type, then label 1 is a label for olive skin tone type, and label 2 is a label for white skin tone type. The following is the GLCM feature extraction result.

Table 1. GLCM Extraction Features

Skin Data	GLCM Extraction Features				
	Angle	Contrast	Corr	Energy	Hmgn
Darkbrown	0°	0.13715	0.97778	0.1764	0.93875
	45°	0.20867	0.96614	0.16398	0.9164
	90°	0.15213	0.97534	0.17349	0.93482
	135°	0.22202	0.96398	0.16216	0.91303
Olive	0°	0.12676	0.98222	0.12811	0.9387
	45°	0.25215	0.96465	0.1071	0.8927
	90°	0.19311	0.973	0.11507	0.91231
	135°	0.23871	0.96653	0.10956	0.8985
White	0°	0.19703	0.96988	0.19408	0.91183
	45°	0.31217	0.95241	0.1819	0.87988
	90°	0.23152	0.96465	0.19254	0.90239
	135°	0.30648	0.95327	0.18486	0.88307

The results show that the distribution of features before normalization has a varied and non-uniform range, while after the Z-score normalization process, the distribution of features becomes more

uniform with a mean of 0 and a standard deviation of 1. This ensures that all features have the same scale, facilitating more accurate analysis in machine learning models.

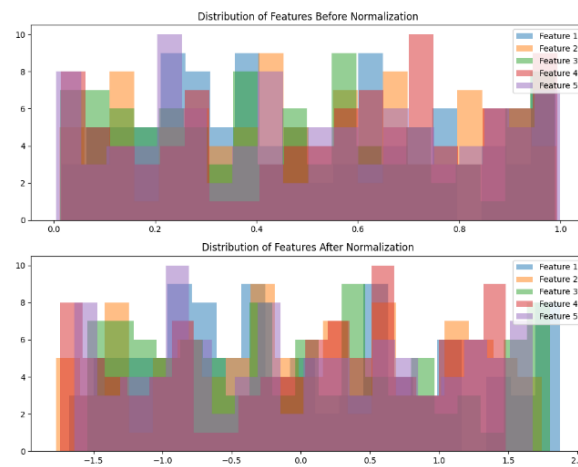


Figure 6. Z-Score Result

3.2. Data Splitting

Next, the data will be cut using cross-validation with a ratio of 80:20, with 80% as training data and 20% as test data because the ratio is stable ratio in cutting data and searching for the best K value from the data to find out how much the best K value will be used in the KNN classification and as a selection of the right K value to avoid overfitting or underfitting. After all, the selection of the best K value in the KNN method is indeed a step in building an effective model and improving the prediction performance of the model by using k_values with a list of K 1-100 from the cross-validation score. From the search for the best K on the entire data, the best K value is obtained using the value of $K = 1$ as the best classification model. The following is the amount of data from each section, which can be seen in Table 2.

Table 2. Data Splitting Results

Data	Total
Training	770
Testing	193
Total	963

3.3. Model Evaluation

The accuracy results of GLCM feature extraction for skin tone image classification using the KNN algorithm show high accuracy, with an accuracy of 75.65% at a $K = 1$ value. Then, the evaluation process is carried out with the confusion matrix tool to see how well the model predicts the test data, as shown in the table below.

Table 3. Confusion Matrix Results

		Prediction			
		Class	Dark Brown	Olive	White
Actual	Dark Brown		84	3	8
	Olive		3	26	9
	White		10	14	36

The table above shows the TP of each class, such as dark brown (TP: 84, FP: 11, FN: 13), olive (TP: 26, FP: 12, FN: 17), and white (TP: 36, FP: 24, FN: 17). After that, the results of accuracy, precision, sensitivity, and f1-score are sought with the calculations below.

$$Acc = \frac{84+26+36}{193} = 0,7565 \quad (11)$$

$$Prec = \frac{\left(\frac{84}{84+11} + \frac{26}{26+12} + \frac{36}{36+24}\right)}{3} = 0,7288 \quad (12)$$

$$Sensitivity = \frac{\left(\frac{84}{84+13} + \frac{26}{26+17} + \frac{36}{36+17}\right)}{3} = 0,7166(13)$$

$$F1 - Score = 2 * \frac{0,7288 * 0,7166}{0,7288 + 0,7166} = 0,7226(14)$$

Based on the above calculations, each multiplied by 100%, the model's accuracy is 75.65%, precision is 72.88%, sensitivity is 71.66%, and f1-score is 72.26%. From these results, it can be concluded that the use of GLCM feature extraction and the KNN algorithm was successfully carried out and obtained a robust model.

4. DISCUSSIONS

Based on the test results, it was found that GLCM feature extraction combined with the KNN algorithm could successfully predict the amount of image data used in 770 test data and 193 training data, with a ratio of 80% test data and 20% training data. The identification of training data was performed using the KNN algorithm with a K value and accuracy of K = 1, achieving 75.65%. The KNN algorithm was found to be successful in identifying skin tone types. This is evident from the high accuracy rate, indicating the potential for this method to be used in real-world applications to help women choose foundation that matches their skin tone. This study also opens up opportunities for future improvements through the use of different algorithms, increased data volume, and combination with other feature extraction methods. The following is a comparison table from previous studies.

Table 4. Comparative Study

Authors	Methods	Accuracy (%)
Khoirunisa & Charibaldi [53]	GLCM + SVM	74.00
Kurniati et al. [54]	GLCM + k-NN (K=100)	67.20
Kirimi [55]	GLCM+SVM	70.00
Proposed Method	GLCM+kNN+Best K Finder	75.65

Table 4 presents a comparative study of classification accuracy between the proposed method and several related works in the same domain. The study by Khoirunisa and Charibaldi applied GLCM feature extraction combined with the Support Vector Machine (SVM) classifier and achieved an accuracy of 74.00%. Kurniati et al., on the other hand, utilized GLCM with k-Nearest Neighbor (k-NN) using a large K value of 100, resulting in a significantly lower accuracy of 67.20%, indicating the negative impact of non-optimal K selection. Kirimi also implemented GLCM with SVM and obtained an accuracy of 70.00%. In comparison, the proposed method, which integrates GLCM feature extraction with k-NN and an optimal K value determined through cross-validation (K = 1), achieved the highest accuracy of 75.65%. These results demonstrate that the proposed method outperforms previous approaches by effectively optimizing the K parameter, thereby improving classification performance.

Moreover, the method maintains computational simplicity, making it well-suited for real-time implementation in mobile applications, particularly in the beauty and e-commerce industries.

5. CONCLUSION

The conclusion does not repeat the sentences in the abstract. This research makes several key contributions to the field of automated skin tone classification and beauty technology applications. First, it demonstrates the effectiveness of combining texture-based GLCM features with KNN classification for skin tone discrimination, achieving 75.65% accuracy on Indonesian skin tone categories. Second, it provides empirical evidence that texture analysis outperforms traditional color-based approaches by 15.65%, establishing a new baseline for foundation shade matching systems. Third, the study contributes to the limited body of knowledge on computer vision applications in the cosmetics industry, particularly for Southeast Asian populations.

The conclusion of this research shows that the use of the Gray Level Co-Occurrence Matrix (GLCM) feature extraction method combined with the K-Nearest Neighbor (KNN) algorithm successfully predicts skin tone data with an accuracy of 75.65% using a value of $K = 1$. This research uses 770 training data and 193 test data, with a ratio of 80% training data and 20% test data. The results show that the combination of GLCM and KNN is effective in identifying skin tone, but the KNN algorithm is not able to handle the case of large datasets with high efficiency. Therefore, there is still room for improvement in the algorithm.

The findings have significant practical implications for the beauty industry. For immediate implementation, we recommend developing mobile applications (Android/iOS) that utilize this GLCM-KNN framework for real-time virtual foundation try-on experiences. Integration with e-commerce platforms can enable automated product recommendation systems, reducing return rates and improving customer satisfaction. Beauty retailers can implement this technology in smart mirrors or kiosks for in-store consultations. Additionally, expanding the dataset to include at least 5,000 diverse samples across different ethnicities and lighting conditions would significantly enhance model generalizability.

Plans include using other machine learning algorithms such as Random Forest, SVM, or CNN to compare results, as well as increasing the amount of image data to enrich the model. Future research should focus on three specific directions: (1) Deep learning approaches using convolutional neural networks with transfer learning from pre-trained models like ResNet or EfficientNet, targeting accuracy improvements to $>85\%$; (2) Hybrid systems combining GLCM texture features with color analysis in multiple color spaces (RGB, HSV, YCbCr) for comprehensive skin tone characterization; (3) Real-time processing optimization through model compression techniques and edge computing implementation, achieving response times <2 seconds for mobile applications. In addition, additional feature extraction methods and testing on various skin tones will be conducted to ensure more accurate and representative results. In addition, the limited data provided by the Kaggle platform can cause the model to be less than optimal because it can happen that the data to be predicted is not in the data, so in the future, it will be combined between the dataset contained in the Kaggle platform and the data that will be taken directly. Long-term research goals should include developing standardized protocols for skin tone measurement in varying lighting conditions, creating multi-modal systems incorporating spectral analysis, and establishing benchmark datasets for reproducible research in cosmetic applications.

The implementation of the research results in real applications is expected to help women choose a foundation that suits their skin tone, providing a practical and efficient solution to the challenge of customizing foundation shades. This work establishes a foundation for the next generation of AI-powered beauty technologies, contributing to the democratization of personalized cosmetic solutions and advancing the intersection of computer vision and consumer applications.

ACKNOWLEDGEMENT

This work was supported by Lambung Mangkurat University and Gazi University, which provided essential resources and assistance.

REFERENCES

- [1] Y.-H. Choi, S. E. Kim, and K.-H. Lee, "Changes in consumers' awareness and interest in cosmetic products during the pandemic," *Fash. Text.*, vol. 9, no. 1, p. 1, Jan. 2022, doi: 10.1186/s40691-021-00271-8.
- [2] Y. Yan, J. Lee, J. Hong, and H. Suk, "Measuring and describing the discoloration of liquid foundation," *Color Res. Appl.*, vol. 46, no. 2, pp. 362–375, Apr. 2021, doi: 10.1002/col.22584.
- [3] L. Liu, R. Huang, S. Yang, and N. Sang, "Detecting skin colors under varying illumination," Nov. 2011, p. 80040N, doi: 10.1117/12.901776.
- [4] R. Kips, L. Tran, E. Malherbe, and M. Perrot, "Beyond Color Correction : Skin Color Estimation In The Wild Through Deep Learning," *Electron. Imaging*, vol. 32, no. 5, pp. 82-1-82–8, Jan. 2020, doi: 10.2352/ISSN.2470-1173.2020.5.MAAP-082.
- [5] C. M. Frisby, "Black and Beautiful: A Content Analysis and Study of Colorism and Strides toward Inclusivity in the Cosmetic Industry," *Adv. Journal. Commun.*, vol. 07, no. 02, pp. 35–54, 2019, doi: 10.4236/ajc.2019.72003.
- [6] E. Markiewicz and O. Idowu, "Melanogenic Difference Consideration in Ethnic Skin Type: A Balance Approach Between Skin Brightening Applications and Beneficial Sun Exposure," *Clin. Cosmet. Investig. Dermatol.*, vol. Volume 13, pp. 215–232, Mar. 2020, doi: 10.2147/CCID.S245043.
- [7] M. Benčević, M. Habijan, I. Galić, D. Babin, and A. Pižurica, "Understanding skin color bias in deep learning-based skin lesion segmentation," *Comput. Methods Programs Biomed.*, vol. 245, p. 108044, Mar. 2024, doi: 10.1016/j.cmpb.2024.108044.
- [8] S. K. Mbatha, M. J. Booysen, and R. P. Theart, "Skin Tone Estimation under Diverse Lighting Conditions," *J. Imaging*, vol. 10, no. 5, p. 109, Apr. 2024, doi: 10.3390/jimaging10050109.
- [9] C. M. Heldreth, E. P. Monk, A. T. Clark, C. Schumann, X. Eyee, and S. Ricco, "Which Skin Tone Measures Are the Most Inclusive? An Investigation of Skin Tone Measures for Artificial Intelligence," *ACM J. Responsible Comput.*, vol. 1, no. 1, pp. 1–21, Mar. 2024, doi: 10.1145/3632120.
- [10] P. Pi and D. Lima, "Gray level co-occurrence matrix and extreme learning machine for Covid-19 diagnosis," *Int. J. Cogn. Comput. Eng.*, vol. 2, pp. 93–103, Jun. 2021, doi: 10.1016/j.ijcce.2021.05.001.
- [11] M. M. D. Oghaz, V. Argyriou, D. Monekosso, and P. Remagnino, "Skin Identification Using Deep Convolutional Neural Network," 2019, pp. 181–193.
- [12] P. Bjerring and P. H. Andersen, "Skin reflectance spectrophotometry," *Photodermatol.*, vol. 4, no. 3, pp. 167–171, Jun. 1987.
- [13] H. K. Jeong, C. Park, R. Henao, and M. Kheterpal, "Deep Learning in Dermatology: A Systematic Review of Current Approaches, Outcomes, and Limitations," *JID Innov.*, vol. 3, no. 1, p. 100150, Jan. 2023, doi: 10.1016/j.xjidi.2022.100150.
- [14] A. Almutairi and R. U. Khan, "Image-Based Classical Features and Machine Learning Analysis of Skin Cancer Instances," *Appl. Sci.*, vol. 13, no. 13, p. 7712, Jun. 2023, doi: 10.3390/app13137712.
- [15] A. Rasak, A. Zubair, O. A. Alo, and A. R. Zubair, "Grey Level Co-occurrence Matrix (GLCM) Based Second Order Statistics for Image Texture Analysis," *Int. J. Sci. Eng. Investig.*, vol. 8, no. 93, p. 93, 2019, [Online]. Available: www.IJSEI.com.
- [16] R. M. Haralick, K. Shanmugam, and I. Dinstein, "Textural Features for Image Classification," *IEEE Trans. Syst. Man. Cybern.*, vol. SMC-3, no. 6, pp. 610–621, Nov. 1973, doi: 10.1109/TSMC.1973.4309314.
- [17] U. Markowska-Kaczmar and A. Slimak, "Creating herd behavior by virtual agents using neural networks," *Procedia Comput. Sci.*, vol. 192, pp. 437–446, 2021, doi: 10.1016/j.procs.2021.08.045.

-
- [18] R. Andrian, M. A. Naufal, B. Hermanto, A. Junaidi, and F. R. Lumbanraja, “k-Nearest Neighbor (k-NN) Classification for Recognition of the Batik Lampung Motifs,” *J. Phys. Conf. Ser.*, vol. 1338, no. 1, p. 012061, Oct. 2019, doi: 10.1088/1742-6596/1338/1/012061.
 - [19] T. Johnson and S. K. Singh, “Fuzzy C strange points clustering algorithm,” in *2016 International Conference on Information Communication and Embedded Systems (ICICES)*, Feb. 2016, pp. 1–5, doi: 10.1109/ICICES.2016.7518929.
 - [20] P. Cunningham and S. J. Delany, “k-Nearest Neighbour Classifiers - A Tutorial,” *ACM Comput. Surv.*, vol. 54, no. 6, pp. 1–25, 2022.
 - [21] S. Khare and K. P. Peeyush, “Automotive drive recorder as black box for low end vehicles,” in *2017 International Conference on Advances in Computing, Communications and Informatics (ICACCI)*, Sep. 2017, pp. 1855–1861, doi: 10.1109/ICACCI.2017.8126115.
 - [22] Kuo Chih Liao, F. J. Richmond, T. Hogen-Esch, L. Marcu, and G. E. Loeb, “Sencil/spl trade/ project: development of a percutaneous optical biosensor,” in *The 26th Annual International Conference of the IEEE Engineering in Medicine and Biology Society*, vol. 3, pp. 2082–2085, doi: 10.1109/IEMBS.2004.1403612.
 - [23] S. A. S. Rosiva, M. Zarlis, and W. Wanayumini, “Identifikasi Citra Daun dengan GLCM (Gray Level Co-Occurrence) dan K-NN (K-Nearest Neighbor),” *MATRIK J. Manajemen, Tek. Inform. dan Rekayasa Komput.*, vol. 21, no. 2, pp. 477–488, Mar. 2022, doi: 10.30812/matrik.v21i2.1572.
 - [24] S. Suharyana, F. Anwar, A. C. Dewi, M. Yunianto, U. Salamah, and R. Chai, “Pneumonia Classification Based on GLCM Features Extraction using K-Nearest Neighbor,” *Indones. J. Appl. Phys.*, vol. 13, no. 2, p. 325, Nov. 2023, doi: 10.13057/ijap.v13i2.77120.
 - [25] D. P. Pamungkas, “Ekstraksi Citra menggunakan Metode GLCM dan KNN untuk Identifikasi Jenis Anggrek (Orchidaceae),” *Innov. Res. Informatics*, vol. 1, no. 2, Oct. 2019, doi: 10.37058/innovatics.v1i2.872.
 - [26] O. Elsherif, “Skin Tone Data Classification,” *Kaggle*, 2022. <https://www.kaggle.com/datasets/omarelsherif010/skin-tone-classification> (accessed Apr. 07, 2024).
 - [27] F. Pérez-García, R. Sparks, and S. Ourselin, “TorchIO: A Python library for efficient loading, preprocessing, augmentation and patch-based sampling of medical images in deep learning,” *Comput. Methods Programs Biomed.*, vol. 208, p. 106236, Sep. 2021, doi: 10.1016/j.cmpb.2021.106236.
 - [28] N. Iqbal, R. Mumtaz, U. Shafi, and S. M. H. Zaidi, “Gray level co-occurrence matrix (GLCM) texture based crop classification using low altitude remote sensing platforms,” *PeerJ Comput. Sci.*, vol. 7, p. e536, May 2021, doi: 10.7717/peerj-cs.536.
 - [29] F. Tampinongkol, “Identifikasi Penyakit Daun Tomat Menggunakan Gray Level Co-occurrence Matrix (GLCM) dan Support Vector Machine (SVM),” *Techno Xplore J. Ilmu Komput. dan Teknol. Inf.*, vol. 8, no. 1, pp. 08–16, Apr. 2023, doi: 10.36805/technoexplo.v8i1.3578.
 - [30] K. Djunaidi, H. Bedi Agtriadi, D. Kuswardani, and Y. S. Purwanto, “Gray level co-occurrence matrix feature extraction and histogram in breast cancer classification with ultrasonographic imagery,” *Indones. J. Electr. Eng. Comput. Sci.*, vol. 22, no. 2, p. 795, May 2021, doi: 10.11591/ijeecs.v22.i2.pp795-800.
 - [31] I. Sutrisno *et al.*, “Design of Pothole Detector Using Gray Level Co-occurrence Matrix (GLCM) And Neural Network (NN),” *IOP Conf. Ser. Mater. Sci. Eng.*, vol. 874, no. 1, p. 012012, Jun. 2020, doi: 10.1088/1757-899X/874/1/012012.
 - [32] P. Kulkarni and D. Madathil, “Echocardiography image segmentation using semi-automatic numerical optimisation method based on wavelet decomposition thresholding,” *Int. J. Imaging Syst. Technol.*, vol. 31, no. 4, pp. 2295–2304, Dec. 2021, doi: 10.1002/ima.22631.
 - [33] P. Schober, E. J. Mascha, and T. R. Vetter, “Statistics From A (Agreement) to Z (z Score): A Guide to Interpreting Common Measures of Association, Agreement, Diagnostic Accuracy, Effect Size, Heterogeneity, and Reliability in Medical Research,” *Anesth. Analg.*, vol. 133, no. 6, pp. 1633–1641, Dec. 2021, doi: 10.1213/ANE.0000000000005773.
 - [34] I. M. Putra, I. Tahyudin, H. A. A. Rozaq, A. Y. Syafa’At, R. Wahyudi, and E. Winarto, “Classification analysis of COVID19 patient data at government hospital of banyumas using
-

- machine learning,” in *2021 2nd International Conference on Smart Computing and Electronic Enterprise: Ubiquitous, Adaptive, and Sustainable Computing Solutions for New Normal, ICSCEE 2021*, Jun. 2021, pp. 271–274, doi: 10.1109/ICSCEE50312.2021.9498020.
- [35] M. Mahmud *et al.*, “Implementation of C5.0 Algorithm using Chi-Square Feature Selection for Early Detection of Hepatitis C Disease,” *J. Electron. Electromed. Eng. Med. Informatics*, vol. 6, no. 2, pp. 116–124, Mar. 2024, doi: 10.35882/jeeemi.v6i2.384.
- [36] C. Andrade, “Z Scores, Standard Scores, and Composite Test Scores Explained,” *Indian J. Psychol. Med.*, vol. 43, no. 6, pp. 555–557, Nov. 2021, doi: 10.1177/02537176211046525.
- [37] D. Könen, D. Schmidt, and C. Spisla, “Finding all minimum cost flows and a faster algorithm for the K best flow problem,” *Discret. Appl. Math.*, vol. 321, pp. 333–349, 2022, doi: 10.1016/j.dam.2022.07.007.
- [38] A. R. Isnain, J. Supriyanto, and M. P. Kharisma, “Implementation of K-Nearest Neighbor (K-NN) Algorithm For Public Sentiment Analysis of Online Learning,” *IJCCS (Indonesian J. Comput. Cybern. Syst.)*, vol. 15, no. 2, p. 121, Apr. 2021, doi: 10.22146/ijccs.65176.
- [39] V. R. Joseph, “Optimal ratio for data splitting,” *Stat. Anal. Data Min. ASA Data Sci. J.*, vol. 15, no. 4, pp. 531–538, Aug. 2022, doi: 10.1002/sam.11583.
- [40] A. A. Kurniawan and M. Mustikasari, “Implementasi Deep Learning Menggunakan Metode CNN dan LSTM untuk Menentukan Berita Palsu dalam Bahasa Indonesia,” *J. Inform. Univ. Pamulang*, vol. 5, no. 4, p. 544, 2021, doi: 10.32493/informatika.v5i4.6760.
- [41] Q. H. Nguyen *et al.*, “Influence of Data Splitting on Performance of Machine Learning Models in Prediction of Shear Strength of Soil,” *Math. Probl. Eng.*, vol. 2021, pp. 1–15, Feb. 2021, doi: 10.1155/2021/4832864.
- [42] G. Valente, A. L. Castellanos, L. Hausfeld, F. De Martino, and E. Formisano, “Cross-validation and permutations in MVPA: Validity of permutation strategies and power of cross-validation schemes,” *Neuroimage*, vol. 238, p. 118145, Sep. 2021, doi: 10.1016/j.neuroimage.2021.118145.
- [43] B. G. Marcot and A. M. Hanea, “What is an optimal value of k in k-fold cross-validation in discrete Bayesian network analysis?,” *Comput. Stat.*, vol. 36, no. 3, pp. 2009–2031, Sep. 2021, doi: 10.1007/s00180-020-00999-9.
- [44] A. X. Wang, S. S. Chukova, and B. P. Nguyen, “Ensemble k-nearest neighbors based on centroid displacement,” *Inf. Sci. (Ny.)*, vol. 629, pp. 313–323, Jun. 2023, doi: 10.1016/j.ins.2023.02.004.
- [45] C. Feng, B. Zhao, X. Zhou, X. Ding, and Z. Shan, “An Enhanced Quantum K-Nearest Neighbor Classification Algorithm Based on Polar Distance,” *Entropy*, vol. 25, no. 1, p. 127, Jan. 2023, doi: 10.3390/e25010127.
- [46] I. M. Karo Karo, A. Khosuri, J. S. I. Septory, and D. P. Supandi, “Pengaruh Metode Pengukuran Jarak pada Algoritma k-NN untuk Klasifikasi Kebakaran Hutan dan Lahan,” *J. Media Inform. Budidharma*, vol. 6, no. 2, p. 1174, Apr. 2022, doi: 10.30865/mib.v6i2.3967.
- [47] M. Suyal and P. Goyal, “A Review on Analysis of K-Nearest Neighbor Classification Machine Learning Algorithms based on Supervised Learning,” *Int. J. Eng. Trends Technol.*, vol. 70, no. 7, pp. 43–48, Jul. 2022, doi: 10.14445/22315381/IJETT-V70I7P205.
- [48] W. Wahyono, I. N. P. Trisna, S. L. Sariwening, M. Fajar, and D. Wijayanto, “Comparison of distance measurement on k-nearest neighbour in textual data classification,” *J. Teknol. dan Sist. Komput.*, vol. 8, no. 1, pp. 54–58, Jan. 2020, doi: 10.14710/jtsiskom.8.1.2020.54-58.
- [49] M. Hasnain, M. F. Pasha, I. Ghani, M. Imran, M. Y. Alzahrani, and R. Budiarto, “Evaluating Trust Prediction and Confusion Matrix Measures for Web Services Ranking,” *IEEE Access*, vol. 8, pp. 90847–90861, 2020, doi: 10.1109/ACCESS.2020.2994222.
- [50] S. Napi’ah, T. H. Saragih, D. T. Nugrahadi, D. Kartini, and F. Abadi, “Implementation of Monarch Butterfly Optimization for Feature Selection in Coronary Artery Disease Classification Using Gradient Boosting Decision Tree,” *J. Electron. Electromed. Eng. Med. Informatics*, vol. 5, no. 4, Oct. 2023, doi: 10.35882/jeeemi.v5i4.331.
- [51] I. Tahyudin, H. A. A. Rozaq, and H. Nambo, “Machine Learning Analysis for Temperature Classification using Bioelectric Potential of Plant,” in *2022 6th International Conference on Information Technology, Information Systems and Electrical Engineering (ICITISEE)*, Dec. 2022, pp. 465–470, doi: 10.1109/ICITISEE57756.2022.10057768.

-
- [52] J. Xu, Y. Zhang, and D. Miao, "Three-way confusion matrix for classification: A measure driven view," *Inf. Sci. (Ny)*., vol. 507, pp. 772–794, Jan. 2020, doi: 10.1016/j.ins.2019.06.064.
 - [53] F. Eka Khoirunisa and N. Charibaldi, "Effect of Using GLCM and LBP+HOG Feature Extraction on SVM Method in Classification of Human Skin Disease Type," *Inf. J. Ilm. Bid. Teknol. Inf. dan Komun.*, vol. 9, no. 2, pp. 145–150, Jul. 2024, doi: 10.25139/inform.v9i2.8275.
 - [54] F. T. Kurniati, D. H. Manongga, I. Sembiring, S. Wijono, and R. R. Huizen, "The object detection model uses combined extraction with KNN and RF classification," *Indones. J. Electr. Eng. Comput. Sci.*, vol. 35, no. 1, p. 436, Jul. 2024, doi: 10.11591/ijeecs.v35.i1.pp436-445.
 - [55] F. Kirimi, "Remote sensing based assessment of land cover and soil moisture in the Kilombero floodplain in Tanzania," Bonn University, 2018.