

ABSTRAK

Nabil Zerma maulana, 2025. **Pengembangan Vision Transformer (ViT) untuk Kurasi Aliran Seni Lukis di Website Galeri Digital**. Tugas Akhir, Program Studi Informatika (S1), Universitas Bhinneka Nusantara, Pembimbing: Subari

Kata kunci: pembelajaran mesin, Vision transformer, klasifikasi gambar

Penelitian ini mengembangkan website galeri digital yang diintegrasikan *Vision Transformer* (ViT) untuk tujuan pengklasifikasian gaya lukis. Tujuan utama dari penelitian ini adalah mengembangkan model vision transformer (ViT) yang mampu mengenali dan mengklasifikasikan gaya lukis secara otomatis. Dataset yang digunakan berasal dari Huggingface dan telah diproses untuk mencakup ribuan citra dengan berbagai gaya. Model dilatih menggunakan vit-base/16 dengan ukuran input 224×224 piksel dan strategi pelatihan yang mencakup augmentasi data serta fine-tuning dari pre-trained weights. Website dibangun menggunakan framework django. Evaluasi model dilakukan menggunakan metrik akurasi, presisi, recall, dan f1-score, dengan hasil akurasi akhir sebesar 75%. Hasil tersebut menunjukkan bahwa ViT memiliki potensi dalam kurasi gaya seni secara otomatis, meskipun masih terdapat ruang untuk peningkatan melalui optimasi hyperparameter, penambahan data, dan penguatan antarmuka pengguna. Penelitian ini diharapkan dapat menjadi kontribusi awal dalam pengembangan kurasi seni berbasis AI dalam ekosistem galeri digital.

ABSTRACT

Nabil Zerma Maulana, 2025. **Development of Vision Transformer (ViT) for Curating Painting Style on Digital Gallery Websites.** Final Project, Study Program Informatics Bachelor's degree, Universitas Bhinneka Nusantara, Advisor 1 : Subari

Keyword: machine learning, Vision transformer, image classification

This study developed a digital gallery website integrated with Vision Transformer (ViT) for the purpose of classifying painting styles. The main objective of this research is to develop a Vision Transformer (ViT) model capable of automatically recognizing and classifying painting styles. The dataset used was sourced from Hugging Face and processed to include thousands of images with various styles. The model was trained using ViT-Base/16 with an input size of 224×224 pixels, employing a training strategy that included data augmentation and fine-tuning from pre-trained weights. The website was built using the Django framework. Model evaluation was carried out using accuracy, precision, recall, and F1-score metrics, with a final accuracy result of 75%. These results indicate that ViT has potential for automatic art style curation, although there is still room for improvement through hyperparameter optimization, data expansion, and user interface enhancement. This research is expected to serve as an initial contribution to the development of AI-based art curation in the digital gallery ecosystem.