

ABSTRAK

Bahasa Jawa merupakan salah satu bahasa daerah dengan jumlah penutur terbanyak di Indonesia, namun penggunaan aksara Jawa (*Hanacaraka*) mengalami penurunan signifikan. Aksara Jawa adalah warisan budaya penting yang perlu dilestarikan melalui inovasi teknologi, seperti sistem pengenalan ucapan (*speech recognition*). Penelitian ini mengembangkan sistem klasifikasi ucapan aksara Jawa dasar menggunakan metode *Mel-Frequency Cepstral Coefficient* (MFCC) untuk ekstraksi ciri dan *Support Vector Machine* (SVM) dengan kernel *Radial Basis Function* (RBF) untuk klasifikasi.

Penelitian ini menggunakan model pengembangan *Waterfall*. Data berupa rekaman audio 20 aksara Jawa dasar berformat .wav dikumpulkan dari sejumlah penutur. Proses *preprocessing* meliputi normalisasi, *pre-emphasis*, *framing*, dan *windowing*. Dataset dibagi dengan rasio 80% data latih dan 20% data uji menggunakan *stratified split*.

Hasil pengujian menunjukkan bahwa penggunaan *raw audio* menghasilkan akurasi yang sangat rendah (12%-17%), sedangkan implementasi MFCC secara signifikan meningkatkan performa model. Akurasi tertinggi pada data latih mencapai 97%, sementara pada data uji mencapai 79%. Hal ini membuktikan bahwa kombinasi MFCC dan SVM efektif dalam mengklasifikasikan ucapan aksara Jawa dasar dan dapat berkontribusi pada upaya pelestarian budaya melalui pendekatan teknologi.

Kata Kunci: Aksara Jawa, Klasifikasi Ucapan, MFCC, SVM, Pengenalan Suara.

ABSTRACT

Javanese is a widely spoken regional language in Indonesia, yet the use of Javanese script (Hanacaraka) has significantly declined. As a vital cultural heritage, it requires preservation through technological innovations such as speech recognition. This study develops a classification system for basic Javanese script speech using Mel-Frequency Cepstral Coefficient (MFCC) for feature extraction and Support Vector Machine (SVM) with a Radial Basis Function (RBF) kernel for classification.

The system development follows the Waterfall model. Audio recordings of 20 basic Javanese syllables in .wav format were collected from several speakers. Preprocessing steps included normalization, pre-emphasis, framing, and windowing. The dataset was divided into 80% training data and 20% test data using a stratified split.

Experimental results show that using raw audio yielded very low accuracy (12%-17%), whereas the implementation of MFCC significantly improved the model's performance. The highest accuracy achieved was 97% for the training data and 79% for the test data. This demonstrates that the MFCC-SVM combination is effective for classifying basic Javanese script speech and contributes to cultural preservation through technological innovation.

Keywords: *Javanese Script, Speech Classification, MFCC, SVM, Speech Recognition.*