

Abstract

Speaker recognition is one of the biometric technologies to recognize humans, such as identification of fingerprints, DNA, and iris because the sound characteristics of each human being are different. Unvoiced or consonant pronunciation is noise at the speaker recognition, so removing unvoiced parts can increase the accuracy. The short time zero-crossing rate (STZCR) method is a method that can detect unvoiced parts so they can be removed before the feature extraction process uses MFCC and GMM. After doing some experiments, found the threshold on STZCR for the VoxForge dataset so that accuracy can be obtained 100 % better than not slicing audio file based on STZCR at 99,91%.

Keywords: speaker recognition, unvoiced, STZCR, threshold, MFCC, GMM