

ABSTRAK

Tujuan penelitian ini terhadap klasifikasi kategori kesejahteraan masyarakat untuk mengevaluasi hasil ketetapan terhadap kategori keluarga yang mampu dan kurang mampu. Dalam penelitian ini data yang di peroleh sebanyak 103 data dengan berbentuk format exel dengan terbagi menjadi 2 kelas yaitu kelas mampu yang di tandai dengan "0" Dan tidak mampu "1" Untuk di klasifikasi dengan menggunakan algoritma *K-Nearest Neighbor*. Kemudian melakukan beberapa tahapan data untuk bisa di klasifikasi ya itu dengan analisis data, seleksi data, transformasi data, implementasi algoritma *K-Nearest Neighbor*, dan evaluasi data untuk menghasilkan klasifikasi kesejahteraan masyarakat. Selanjutnya melakukan pembagian dataset dibagi menjadi data uji sebanyak 82 data, dan data latih sebanyak 21 data dengan menggunakan rasio 0.2 yang artinya 80% sebagai data latih dan 20% sebagai data uji, tahapan selanjutnya melakukan evaluasi dalam sebuah pengujian model klasifikasi algoritma *K-Nearest Neighbor* dengan *confusion matrix*, dari pengujian tersebut mendapatkan hasil nilai *accuracy* yang cukup baik dengan nilai *accuracy* 80%, *Precision* 75%, dan *Recall* 90%.

Kata kunci: Klasifikasi, Kesejahteraan Masyarakat, *K-Nearest Neighbor*, *Phyton*, *Confusion matrix*.



ABSTRACT

The purpose of this research on the classification of community welfare categories is to evaluate the results of the determination of the categories of well-off and poor families. In this study the data obtained was 103 data in the form of excel format with divided into 2 classes, namely the capable class marked with "0" and incapable "1" to be classified using the K-Nearest Neighbor algorithm. Then perform several stages of data to be classified, namely data analysis, data selection, data transformation, implementation of the K-Nearest Neighbor algorithm, and data evaluation to produce a classification of community welfare. Furthermore, the dataset is divided into test data as much as 82 data, and training data as much as 21 data using a ratio of 0.2 which means 80% as training data and 20% as test data, the next step is to evaluate in a test of the K-Nearest Neighbor algorithm classification model with a confusion matrix, from the test the results get a pretty good accuracy value with an accuracy value of 80%, Precision 75%, and Recall 90%.

Keywords: Classification, Community Welfare, K-Nearest Neighbor, Python, Confusion matrix.

