

Automatic Recognition of the NIK in Electronic KTP

1st Sekar Mirah¹, 2nd Anastasia Rita Widiarti²
{sekarmirah012@gmail.com¹, rita_widiarti@usd.ac.id²}

Department of Informatics, Faculty of Science and Technology,
Sanata Dharma University, Indonesia¹²

Abstract. Automatic recognition of NIK is a process carried out to perform numeric character recognition stages. The purpose of automatic recognition of NIK is to produce a text that contains images by processing E-KTP images. The preprocessing stage has 4 steps, namely resizing, grayscaling, binarization and segmentation. The segmentation method which is the projection profile. The extraction phase uses invariant moments of HU and Intensity of Character (IoC). The accuracy of our results obtained for the segmentation process is 98.57% of 70 E-KTP data. The accuracy of our results obtained to reach the peak has a value of 1-fold of 86.96%, a 2-fold of 91.30% and a 3-fold of 78.26%.

Keywords: Segmentation, invariant moment, intensity of character.

1 Introduction

A proof of identification that must be carried out by Indonesian citizens for 17 years-old or above is an electronic KTP (E-KTP). When using public transport such as trains or planes, you will use identification. At present there is no tool that can check or automate public transportation user identity numbers. The distinctive feature of E-KTP is NIK, each E-KTP has a different NIK for each person.

Character recognition is a way that is done by processing an image into a number or letter. Character recognition is done using the pattern recognition method. Pattern recognition is a machine of an automatic recognition machine, description, classification, and pattern grouping which is an important problem for engineering and disciplines such as biology, psychology, medicine, marketing, computer vision, artificial intelligence, and remote sensing. The aim of pattern recognition is to perform supervised and unsupervised data classification [1]. This research will conduct automatic recognition of NIK E-KTP by using pattern recognition models namely preprocessing, feature extraction, classification and evaluation.

Rahul and Kumar [2] conducted research on the segmentation of text lines in handwriting. The method used in this research is projection profile. This research has carried out several assumptions used to build an algorithm. The assumptions used include determining the minimum height, average height, and maximum height of text lines. This assumption will be used to determine the text line; if the text line is less than the minimum height then it will be combined with the previous line or after that depending on the line conditions. In addition, this assumption is also used to determine the text line that must be broken with the condition that the height of the text line is greater than the average height summed to the minimum height.

Widiarti, et al. [3] conducted research on preprocessing for Javanese manuscripts. Preprocessing steps carried out are binarization, noise reduction, line segmentation, and character segmentation. The method used for the segmentation process is the projection profile. In this research, there are several assumptions for the segmentation process; the assumption that has been determined is used to check a normal and abnormal line and character. Lines and characters are said to be normal if they are in a range of values, the range is derived from the average value of the line and character and the standard deviation of the line and character. The normal limit of a line or character is if it is located in the range of the average value minus the standard deviation and the average value is summed with the standard deviation value.

Belagali and Angadi [4] research carried out on OCR (Optical Character Recognition) for Canadian character handwriting. The process for developing an OCR is preprocessing, feature extraction, and classification. The segmentation method used is connected component labeling. Whereas for feature extraction using Hu moment invariant method, horizontal, and vertical projection from zone division. The classification process is used by using PNN (probabilistic of neural networks). The zoning division for feature extraction is conducted for each character. The zone division is conducted by dividing it into 3 sections horizontally.

Sridevi and Subashini [5] conducted research on feature extraction using methods using momentary basic in Tamil handwriting. The moment method used is zernike, invariant moment, affine moment. In addition, the feature extraction method used is regional characteristics. The classification phase uses the PNN method. This research compares the accuracy result of each method and combines several methods.

2 Research Method

E-KTP has the same characteristics for each region that has the same order of contents. The NIK on the E-KTP is located on the 3rd row, in addition to the 3rd row there are letters and numbers[6]. To do character recognition, NIK requires a model to compile a system. The pattern recognition model for this research is described in **Figure 1**.

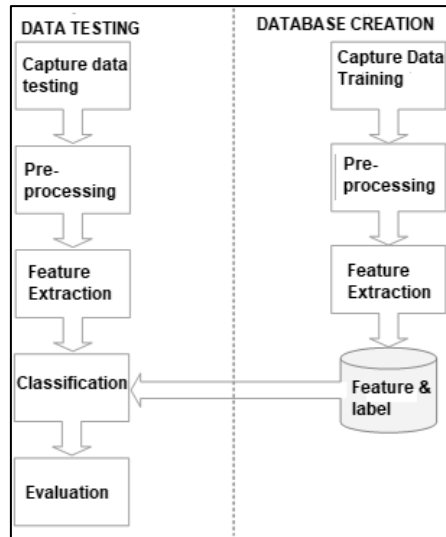


Fig. 1.Model of NIK-recognition.

The pattern recognition model is divided into 2 parts, namely data testing and database creation. Phase 1 to phase 3 has the same process for the data testing and database creation section. The first stage is data sharing between training data and testing data. Distribution of data using 3-fold cross validation with the composition of training data as much as 2/3 of 70 E-KTP and testing data as much as 1/3 of 70 E-KTP. In stage 4, the creation of the database will be carried out storage features of the feature extraction result. The features generated from each character are 30 features. Whereas in stage 4 data testing is a classification process, the classification process is carried out using the template matching method. The last stage of data testing is evaluation, evaluation is conducted to calculate the accuracy the value of segmentation process, calculate the value of recognition accuracy of NIK for each E-KTP and calculate the total accuracy value for the introduction of NIK for all E-KTP data. The calculation of the total accuracy value is conducted for each fold that has been determined.

The introduction of numeric characters from the NIK E-KTP will go through the preprocessing stage; this stage has the input in the form of the original E-KTP image taken using a smartphone.

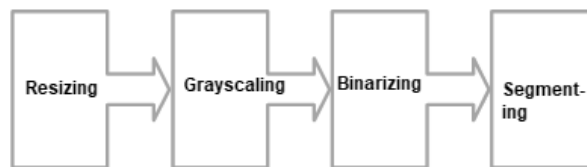


Fig. 2.Block Diagram Preprocessing.

The next step is the preprocessing stage. The steps for preprocessing are illustrated in **Figure 2**. The first step is resizing, this stage is conducted because the size of the E-KTP image is too large so the process carried out by Matlab is too heavy. The initial size of the image is 2448 x 3264 then resizing with a size of 300 x 500. This size is used because the E-KTP image is still clearly visible and considered to be right. The next stage is grayscaleing, at

this stage using the Matlab function, `rgb2gray`. The result of the grayscale image will be the input for the binarization stage; binarization is conducted using the Matlab function that is `imbinarize`. In the binarization function, a value is determined for the threshold. The value used is 0.3. The last preprocessing stage is segmentation, but before doing segmentation additional processes are needed to delete photos on E-KTP. Photos on E-KTP need to be deleted because it will affect the success of the segmentation stage. Deleting photos is conducted by changing the photo object into a background color. The size of the photo on the E-KTP is almost the same size, obtained by the size of the object that has a length of more than 20 and a width of more than 12 is considered as a photo object. If the object meets this size, it will be converted to 0.

The segmentation stage is divided into 2 namely vertical projection and horizontal projection. Vertical projection will divide object lines on E-KTP, while horizontal projection will divide characters from vertical projection results.

Figure 3 is an illustration of the curve for vertical projection. The x axis is the number of white pixels or object pixels and the y axis is the row. The curve is obtained using equation (1). Vertical projections will add each row to all of its columns.

$$P_{ver}[k] = \sum_{j=1}^M I[b_k, k] \quad (1)$$

Based on **Figure 3**, it can be seen that each line can already be seen its limits. On the E-KTP the location of the NIK is on the 3rd row but it is uncertain that the NIK is located on the 3rd row if there is noise in the previous row. To obtain the NIK line, tolerance is required for the number of lines to be checked. This research determines the tolerance of line 1 to line 5. If seen from Figure 3, each row has a different size, so this research calculates the average height of NIK and NIK standard deviation for NIK selection. From the data obtained, the average height is 14.11 and the standard deviation is 1.79. The normal height of NIK if it has a greater range than the average height is reduced by the standard deviation and the height of the NIK is less than equal to 22. The value of 22 is obtained from the maximum value of the height of the NIK obtained.

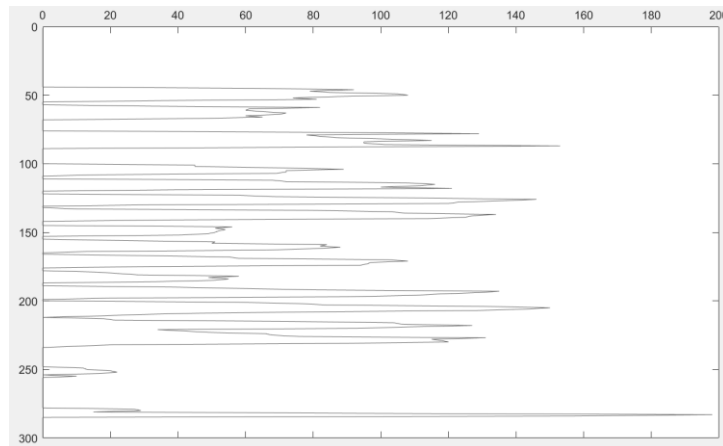


Fig. 3. Vertical Curves of E-KTP Projection.

The results of vertical segmentation will be the input for character segmentation using horizontal projection. After the input results are carried out horizontal projection and implemented in graphical form, the results are as shown in **Figure 4**.

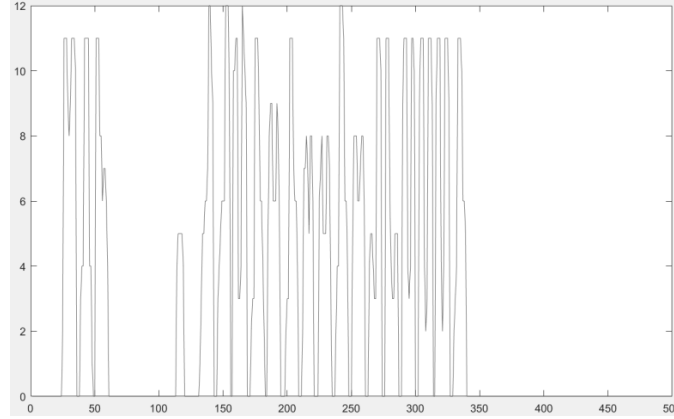


Fig. 4.Horizontal Curve of NIK E-KTP Projection.

Figure 4 is an illustration of the curve for horizontal projections. The x axis is the column and the y axis is the number of white pixels or objects. The curve is obtained using equation (2). Horizontal projections will add each column to all rows.

$$P_{hor}[b] = \sum_{j=1}^N I[b, k_j] \quad (2)$$

The number in NIK is not located in the first character but in characters 1 to 4 which is the writing of NIK and a colon. The purpose of this research is to conduct an introduction to NIK. Therefore, a selection is needed to get numeric characters from NIK. The number of characters in NIK is 16 characters, if the character starts from the 5th character, it will finish at the 21st character. However, it happens if the E-KTP is really clean and there is not much noise on the NIK. After doing research, tolerance is needed because some E-KTPs have noise in the NIK so that it can cause failure to get all NIK numbers. After the experiment, tolerance was used of 2 characters to overcome noise so that the character will be checked up to the 23rd character. In addition, characteristics for the character are needed, so it is also necessary to calculate the average and standard deviation width of the character. After calculated from several data obtained an average value of 11.2 and a standard deviation of 0.89. Besides that the maximum value of the character width is 16. A character is said to be normal if the character has a width less than the same as the average value minus the standard deviation and the character width is not more than the maximum value. If the character is less than the normal width, the character is considered not an object or noise. If the character is more than the maximum value then the character is not normal, it could be said that the object consists of 2 or more characters.

The next stage is the feature extraction stage; this stage will perform the calculations to produce a characteristic of each character. The method that will be used is invariant moment HU and Intensity of Character.

Before calculating the moments of each character, each object will be divided into 3 zones as shown in **Figure 5**. The zoning division is conducted horizontally, then each zone will be

calculated its features using the invariant moment HU method. Feature calculations use invariant moments HU using equation (7) to equation (13).

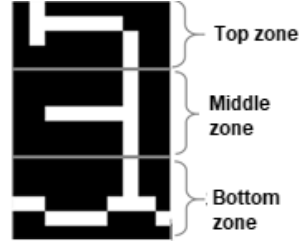


Fig. 5.Zone Division of Invariant Moments HU.

Invariant moments HU are obtained by calculating the value of spatial moments first. The zoning division as shown in **Figure 5** is carried out according to the reference used [4]. Spatial moments can be calculated by using equation (3). The spatial moment will add the value of the column that has been raised with the order value I and multiplied by the row raised by the value of order j .

$$M_{ij} = \sum_{x=1}^M \sum_{y=1}^N x^i y^j I(x, y) \quad (3)$$

In addition, the invariant moment HU uses the central moment to obtain its value. The center moment can be calculated using equation (4). Calculate the center of mass by means of equation (5).

$$\mu_{ij} = \sum_{x=1}^M \sum_{y=1}^N (x - \bar{x})^i (y - \bar{y})^j I(x, y) \quad (4)$$

$$\bar{x} = \frac{M_{10}}{M_{00}}, \bar{y} = \frac{M_{01}}{M_{00}} \quad (5)$$

In order for the central moment to overcome translation, scale and rotation, the central moment needs to be normalized. Normalization can be conducted in a way like equation (6).

$$\eta_{ij} = \frac{\mu_{pq}}{\mu_{00}^\gamma}, \gamma = \frac{i+j+2}{2} \quad (6)$$

Calculating the moment of normalization center is conducted by calculating the order moment value divided by the moment value with order 00.

The characteristic used is invariant moment HU, this method uses calculation with center moment that has been normalized but with various order values. Invariant moment HU has 7 moment values produced.

$$\emptyset_1 = \eta_{20} + \eta_{02} \quad (7)$$

$$\emptyset_2 = (\eta_{20} + \eta_{02})^2 + (4\eta_{11})^2 \quad (8)$$

$$\emptyset_3 = (\eta_{30} - 3\eta_{12})^2 + (3\eta_{21} - \eta_{03})^2 \quad (9)$$

$$\emptyset_4 = (\eta_{30} + \eta_{12})^2 + (\eta_{21} + \eta_{03})^2 \quad (10)$$

$$\emptyset_5 = (\eta_{30} - 3\eta_{12})(\eta_{30} + \eta_{12})\{(\eta_{30} + \eta_{12})^2 - 3(\eta_{21} + \eta_{03})^2\} + (3\eta_{21} - \eta_{03})(\eta_{21} + \eta_{03})\{3(\eta_{30} + \eta_{12})^2 - (\eta_{21} + \eta_{03})^2\} \quad (11)$$

$$\emptyset_6 = (\eta_{20} - \eta_{02})\{(\eta_{30} + \eta_{12})^2 - (\eta_{21} + \eta_{03})^2\} + 4\eta_{11}(\eta_{30} + \eta_{12})(\eta_{21} + \eta_{03}) \quad (12)$$

$$\emptyset_7 = (3\eta_{21} - \eta_{30})(\eta_{30} + \eta_{12})\{(\eta_{30} + \eta_{12})^2 - 3(\eta_{21} + \eta_{30})^2\} + (3\eta_{21} - \eta_{03})(3\eta_{21} + \eta_{03})\{3(\eta_{30} + \eta_{12})^2 - (\eta_{21} + \eta_{03})^2\} \quad (13)$$

In addition to the feature extraction method used is Intensity of Character. This method is also carried out by dividing the area as shown in Figure 6.

The zoning division is conducted vertically, in the previous step the zone has been divided horizontally. From each horizontal zone, it will be divided as shown in **Figure 6**. Each zone will count its characteristics, the resulting characteristics are 3. Intensity of Character method will calculate the number of white pixels or pixels that form objects.

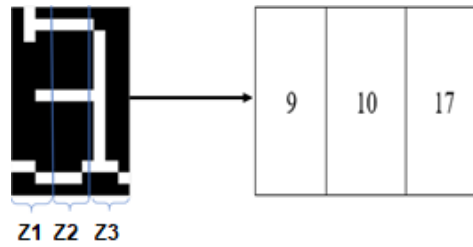


Fig. 6. Zones Distribution of Intensity of Character.

In the data testing, the next stage is classification. The classification method used is template matching. This method will calculate the distance from each feature to all databases. The calculation of distance is using the Euclidean distance method. To calculate the distance, it can be seen in the equation (14).

$$d_{Euc} = \sqrt{\sum_{i=1}^d |P_i - Q_i|^2} \quad (14)$$

Tests are carried out to calculate the success rate of segmentation and the introduction of NIK E-KTP. Segmentation accuracy is carried out for each E-KTP and in total such as equation (15) and (16), while the introduction is conducted to calculate the accuracy of data testing for each fold and the accuracy of recognition of each E-KTP such as equation (17) and (18).

$$segmentation \text{ accuracy of each } E - KTP = \frac{\sum right \text{ segmentation}}{16} \times 100\% \quad (15)$$

$$total \text{ segmentation accuracy} = \frac{\sum E-KTP \text{ with } 100\% \text{ accuracy}}{70} \times 100\% \quad (16)$$

$$recognition \text{ accuracy of each } E - KTP = \frac{\sum right \text{ character}}{16} \times 100\% \quad (17)$$

$$fold \text{ accuration} = \frac{\sum \text{Objek that are correctly recognized}}{\sum 23} \times 100\% \quad (18)$$

3 Result and Analysis

This system will take the NIK E-KTP line for the recognition process, line cutting is conducted with vertical projection. **Figure 7** is the result of a line segmentation process. If the NIK meets the specified conditions, the result of the line segmentation is the same as **Figure 7**. After obtain the NIK line, the result will be the input for the character segmentation stage. **Figure 8** is the result of horizontal projection; all numeric characters are well segmented.



Fig. 7. Vertical Projection of NIK E-KTP.

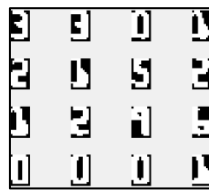


Fig. 8. Horizontal Results of NIK E-KTP Projection.

After successful character segmentation, the next stage is feature extraction. **Figure 9** is an example of a feature extraction that results from a number character. Columns 1 through 21 are the result of invariant moment values and columns 22 to 30 are the results of Intensity of Character. .

	1	2	3	4	5	6	7	8	9	10
1	0.2826	0.1257	0.0056	0.0013	2.2397e-06	9.3329e-05	1.6551e-06	0.1862	0.0091	0.0027
	11	12	13	14	15	16	17	18	19	20
1	8.6078e-05	5.3227e-09	-1.1183e-06	1.2819e-08	0.2634	0.0465	0.0021	3.1494e-04	8.7390e-08	1.7199e-06
	21	22	23	24	25	26	27	28	29	30
1	-4.3767e-09	4	12	7	1	15	12	7	13	6

Fig. 9. Result of Feature Extraction.

The results of the classification stage are described as in **Figure 10**. Column 1 is the result of the value with the smallest distance and the second column is the label of the character.

	1	2		
1	2.8288	"3"	9	1.7322
2	2.4496	"3"	10	2.4500
3	1.7323	"0"	11	4.2427
4	3.0004	"1"	12	2.4498
5	1.0001	"2"	13	2.8284
6	3.3178	"1"	14	1.0000
7	3.4643	"5"	15	1.7322
8	4.6916	"2"	16	2.2362
				"1"

Fig. 10. Result of Classification.

The success of a system is assessed based on the results of the accuracy obtained. The results of data testing will be displayed in graphical form. The test results are in the form of percentage of segmentation and introduction of NIK.

Based on the experiment of segmentation testing using all data, the experimental results are described in a graph as in **Figure 11**, there is 1 data that does not have an accuracy of 100%.

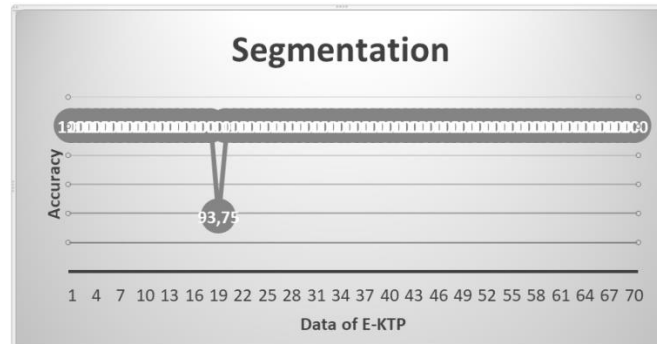


Fig. 11. Segmentation Accuracy Diagram.

This data only has an accuracy of 93.75% with error of 1 character segmentation. The errors due to poor photos and poor lighting eliminate characters. So, to calculate the accuracy of the segmentation process, it is calculated by equation (19).

$$\frac{\text{Number of right data}}{\text{Number of Total Data}} \times 100\% = \frac{69}{70} \times 100\% = 98.57\% \quad (19)$$

Based on the experiments of NIK introduction using data testing one fold, the results are illustrated in a graph like **Figure 12**. The graph illustrates from 23 data there are 3 data that are not successful through the introduction stage. These three data experience character recognition errors caused by noise, imperfect character numbers or parts of missing numbers.

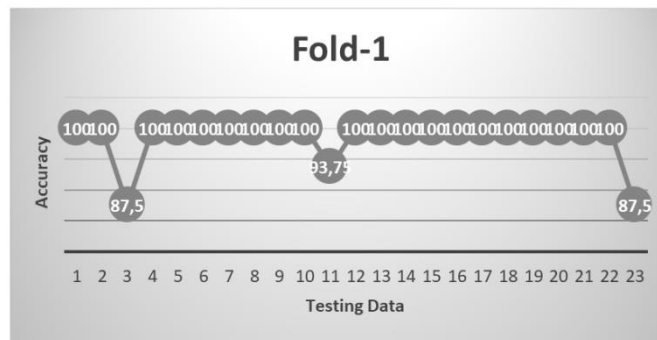


Fig. 12. Diagram of accuracy of recognition with one fold data.

There are 2 E-KTP that incorrectly recognize 2 numeric characters; in addition there is 1 E-KTP that incorrectly identifies 1 numeric character. Each accuracy of E-KTP is calculated by calculating the truth of character recognition for 16 NIK numbers. In addition, it is necessary to calculate the total accuracy of the system by equation (20).

$$\frac{\text{Number of right data}}{\text{Number of Total Data}} \times 100\% = \frac{20}{23} \times 100\% = 86.96\% \quad (20)$$

Based on the experiment of NIK introduction using data testing fold 2, the results are illustrated in a graph like **Figure 13**. The graph illustrates there are 2 data that experienced an error of recognition. Both of these data experience character recognition errors caused by noise, imperfect character numbers or parts of missing numbers. There are 2 E-KTP that incorrectly recognize 2 numeric characters. Each accuracy of E-KTP is calculated by calculating the truth of character recognition for 16 NIK numbers. In addition, the total accuracy of the system needs to be calculated by equation (21).

$$\frac{\text{Number of right data}}{\text{Number of Total Data}} \times 100\% = \frac{21}{23} \times 100\% = 91.30\% \quad (21)$$

Based on the experiment of NIK introduction using data testing fold 3, the results are illustrated in a graph like **Figure 14**. The graph illustrates there are 5 data that experienced an error of recognition. These five data experience character recognition errors caused by noise, imperfect numerical characters or parts of missing numbers and character similarities. Each accuracy of E-KTP is calculated by calculating the truth of character recognition for 16 NIK numbers. In addition, the total accuracy of the system needs to be calculated by equation (22).

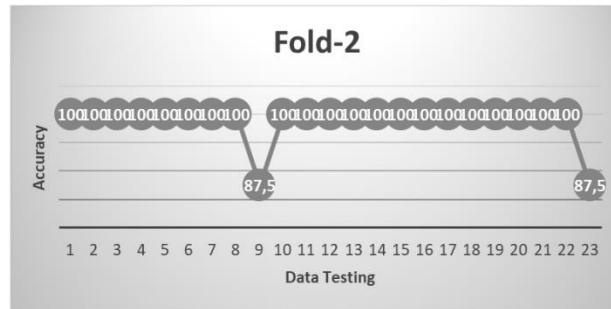


Fig. 13.Diagram of accuracy of recognition with 2-fold data.

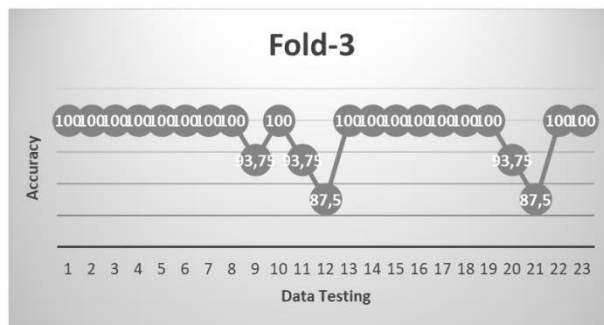


Fig. 14.Diagram of accuracy of recognition with 3-fold data.

$$\frac{\text{Number of right data}}{\text{Number of Total Data}} \times 100\% = \frac{18}{23} \times 100\% = 78.26\% \quad (22)$$

Conclusion

The introduction of NIK on E-KTP has stages starting from preprocessing, feature extraction, and recognition or classification. The sub process of preprocessing is segmentation; this stage uses the projection profile method. Projection profile segmentation has weakness for data that is skewed and has a lot of noise so that characters blend together. In the extraction process the characteristics are very influential for number recognition. Feature extraction method uses invariant moments HU and Intensity of Character has a weakness to classify numbers that are similar in shape or the number that has almost the same pixels. After going through the introduction phase, the accuracy of the system will be calculated. Accuracy obtained from the introduction of NIK on E-KTP with cementation using projection profile and feature extraction methods using moment invariant HU and Intensity of Character results in fairly good accuracy with accuracy of fold 1 of 86.96%, fold 2 of 91.30% and fold 3 amounting to 78.26%.

References

- [1] A. . Jain, R. P. . Duin, and J. Mao, "Statistical Pattern Recognition. A Review. Journal IEEE Transactions on Pattern Analysis and Machine Intelligence," pp. 4 – 37, 2000.
- [2] R. Garg and N. . Garg, .: "An Algorithm for Text Line Segmentation in Handwritten Skewed and Overlapped Devanagari Script. International Journal of Emerging Technology and Advance Engineering," pp. 114–118, 2014.
- [3] A. . Widiarti and Harjoko.A, "Marsono, and Hartati.S.: Preprocessing Model of Manuscript in Javanese Characters. Journal of Signal and Information Processing," p. 112, 2012.
- [4] N. Belagali and A. . Angadi, "OCR for Handwritten Kannada Language Script. International Journal of Recent Trends in Engeneering and Research," pp. 190–197, 2016.
- [5] N. Sridevi and P. Subashini, "Moment Based Feature Extraction for Classification of Handwritten Ancient Tamil Scripts. International Journal of Emerging trends in Engeneering and Development.," pp. 106–115, 2012.
- [6] K. Abdul and S. Adhi, *Teori dan Aplikasi Pengolahan Citra. Andi Offset*. Yogyakarta, 2013.