

ICFHR 2018 Competition On Document Image Analysis Tasks for Southeast Asian Palm Leaf Manuscripts

Made Windu Antara Kesiman^{1,2*}, Dona Valy^{3,4}, Jean-Christophe Burie¹, Erick Paulus⁵,
Mira Suryani⁵, Setiawan Hadi⁵, Michel Verleysen³, Sophea Chhun⁴, and Jean-Marc Ogier¹

¹*Laboratoire Informatique Image Interaction (L3i), Université de La Rochelle, La Rochelle, France*

²*Laboratory of Cultural Informatics (LCI), Universitas Pendidikan Ganesha, Singaraja, Bali, Indonesia*

³*Institute of Information and Communication Technologies, Electronic, and Applied Mathematics (ICTEAM),
Université Catholique de Louvain, Louvain, Belgium*

⁴*Department of Information and Communication Engineering, Institute of Technology of Cambodia, Phnom Penh, Cambodia*

⁵*Department of Computer Science, Universitas Padjadjaran, Bandung, Indonesia*

^{*}) made_windu_antara.kesiman@univ-lr.fr

Abstract—This paper presents the results of the Competition on Document Image Analysis Tasks for Southeast Asian Palm Leaf Manuscripts that was organized in the context of the 16th International Conference on Frontiers in Handwriting Recognition (ICFHR-2018). For this competition, three different corpus of palm leaf manuscripts written in three different scripts and languages (Balinese, Sundanese and Khmer) are used. Four Document Image Analysis (DIA) tasks are proposed as the challenges in this competition: binarization, text line segmentation, isolated character/glyph recognition, and word transliteration. The results of this competition will be very useful in benchmarking analysis for the collection of palm leaf manuscripts, accelerating, evaluating and improving the performance of existing DIA system for a new type of document collection. This paper describes the competition details including the dataset, the evaluation measures used, a short description of each participant as well as the performance of the all submitted methods.

Index Terms—palm leaf manuscript, binarization, text line segmentation, isolated character/glyph recognition, word transliteration, performance evaluation, competition

I. INTRODUCTION

Palm leaves were historically used as writing supports in manuscripts from Southeast Asia. The existence of ancient palm leaf manuscripts in Southeast Asia is very important both in term of quantity and variety of historical contents. This competition is part of an effort to explore Document Image Analysis (DIA) research for palm leaf manuscripts collection as the heritage documents from Southeast Asia. This collection offers a new challenge for DIA researchers because it uses palm leaf as writing media and also with a language and script that have never been analyzed before. The challenges range widely from binarization [1]–[3], text line segmentation [4], [5], character and text recognition tasks [3], [6], [7] to the word spotting methods [8]. The technical challenges can be seen under two points of view: the physical condition of the palm leaf manuscript which strongly influences the quality of captured document images and the

complexity of the script. For this competition, three different corpora of palm leaf manuscripts written in three different scripts and languages have been collected from various locations in Indonesia (Balinese and Sundanese) and Cambodia (Khmer) (Figure 1). The sample images of Balinese palm leaf manuscripts have been collected from 5 different locations (regions) in Bali [9]. For Khmer palm leaf manuscripts, a standard digitization campaign was conducted in order to capture and collect palm leaf manuscript images found in Buddhist temples in different locations throughout Cambodia: Phnom Penh, Kandal, and Siem Reap [10]. The collection of Sundanese palm leaf manuscripts comes from Situs Kabuyutan Ciburuy, Garut, West Java, Indonesia [11]. Four DIA tasks are proposed as the challenges in this competition: binarization, text line segmentation, isolated character/glyph recognition, and word transliteration. Each challenge has a single track with a mixed collection of palm leaf manuscripts from Bali, Cambodia and Sunda and/or some more tracks for each one dedicated to a collection in a specific language. The results of this competition will be very useful in benchmarking analysis for the collection of palm leaf manuscripts, accelerating, evaluating and improving the performance of existing DIA system for a new type of document collection. This paper is organized as follow: Section II gives a brief description about the participants. Section III–VI present the detail description of each challenge. Finally, some conclusions are given in Section VII.

II. THE PARTICIPANTS

Twenty-two research groups registered to this competition: sixteen groups for Challenge A, nine groups for Challenge B, seven groups for Challenge C, and five groups for Challenge D. But only eight research groups accomplished the task and submitted their results: four groups for Challenge A, one group for Challenge B, two groups for Challenge C, and



Fig. 1. Sample images from each corpus

two groups for Challenge D. Brief introduction of the eight research groups who submitted their results are as follows:

- **Group 3** (Challenge D): Submitted by Kartik Dutta, from IIIT-Hyderabad, India.
- **Group 12** (Challenge A): Submitted by XIONG Wei, XIONG Zijie, JIA Xiuhong, and LI Min, from School of Electrical and Electronic Engineering, Hubei University of Technology, Wuhan, China.
- **Group 13** (Challenge D): Submitted by Jianshu Zhang, Jun Du, and Lirong Dai, from National Engineering Laboratory for Speech and Language Information Processing, University of Science and Technology of China, China.
- **Group 17** (Challenge A and B): Submitted by Deepak Kumar, from Department of Electronics & Communication Engineering, Dayananda Sagar Academy of Technology and Management (DSATM), Bengaluru, India.
- **Group 18** (Challenge A): Submitted by Khairun Sadami, Khairunnisa, Putri Afrah, Merita, and Twk Mohd Iqbal, from Multimedia Signal Processing Research Group, Syiah Kuala University, Indonesia.
- **Group 20** (Challenge C): Submitted by Zi-Rui Wang, Jun Du and Wen-Chao Wang, from University of Science and Technology of China, China.
- **Group 21** (Challenge C): Submitted by Cristinel Codrescu, University of Salzburg, Austria.
- **Group 22** (Challenge A): Submitted by Chandranath Adak, School of Software, University of Technology Sydney, Australia.

Brief descriptions of their methods are given in each section of the challenge.

III. CHALLENGE A: BINARIZATION

A. Description of the Challenge

Binarization is widely applied as the first pre-processing step in image document analysis [12]. Binarization is a common starting point for document image analysis. Many binarization methods have been reported. These methods have been tested and evaluated on different types of document collections. A binarization method that performs well for one document collection, may not necessarily be applied to another document collection with the same performance [12]. For this reason, there is always a need to evaluate the new binarization methods for a new document collection that has different characteristics, for example the historical archive documents. The objective of this challenge is to propose a binarization

TABLE I
PALM LEAF MANUSCRIPT DATASETS FOR BINARIZATION TASK

Manuscripts	Train	Test	Dataset
Balinese	50 pages	50 pages	Extracted from AMADL_LontarSet ¹ [3], [9]
Khmer	23 pages	23 pages	Extracted from EFEO ² [13]
Sundanese	31 pages	30 pages	Extracted from Sunda Dataset ICDAR2017 [11]

technique specifically suited for ancient palm leaf documents with degradations and noises.

B. Dataset

The palm leaf manuscript datasets for binarization task are presented in Table I. One ground truth binarized image is provided for each image. The ground truth creation method can be found in [9], [11], [13]. The training set with the ground truth is provided for supervised learning methods. This challenge has only one single track, binarization for a mixed collection of palm leaf manuscripts from Bali, Cambodia and Sunda.

C. Performance Evaluation

Three metrics of binarization evaluation proposed in the DIBCO 2009 contest [14] are used in the evaluation of the binarization task challenge. Those three metrics are F-Measure (FM), Peak SNR (PSNR), and Negative Rate Metric (NRM) [2]. A higher F-Measure and PSNR and a lower NRM indicate a better match.

D. Participant Methods

G12: The proposed method consists of three main stages. Mathematical morphology (MM) is first carried out to compensate the document background with a disk-shaped structuring element, whose the size is determined by stroke width transform (SWT) [15]. If the text pixels are darker than the background pixels, the morphological bottom-hat transform is used, but if the text pixels are brighter than the background, the morphological top-hat transform is applied. A vote is applied to detect whether a dark text is present on a bright background or vice versa based on the SWT with minimum entropy. Secondly, the Howe's binarization method [16] is performed on the compensated document images to further segment the foreground and background pixels. Once the Laplacian energy-based segmentation is derived, a final post-processing stage allows to produce better results.

G17: This method uses difference of Gaussian and non-linear enhancement. The red colored gray values of Palm Leaf RGB image is selected for binarization. The image is enhanced by passing through difference of Gaussian operators. Another round of enhancement is performed using a non-linear function. The non-linearly enhanced image is segmented using a simple threshold value of 0.9. The strokes in the manuscripts are segmented using a simple threshold. The boundary components and noisy elements are removed from the image.

TABLE II
EVALUATION RESULTS OF CHALLENGE A

Group	FM	NRM	PSNR
G12	49.60	0.16	23.43
G17	58.87	0.17	28.71
G18	37.87	0.31	22.87
G22	31.17	0.35	19.10

TABLE III
EVALUATION RESULTS OF CHALLENGE A PER COLLECTION

Group	Bali			Khmer			Sunda		
	FM	NRM	PSNR	FM	NRM	PSNR	FM	NRM	PSNR
G12	52.90	0.16	26.53	44.38	0.15	20.73	48.09	0.17	20.32
G17	56.45	0.19	29.53	66.94	0.11	30.70	56.72	0.20	25.82
G18	47.54	0.22	26.16	1.03	0.59	13.99	49.99	0.24	24.20
G22	40.87	0.30	25.95	3.82	0.62	9.88	35.96	0.22	14.74
G2-2016	68.76	0.13	33.39	-	-	-	-	-	-

G18: The method is an improved NICK (iNICK) binarization method [17]. This method was proposed to set k value of NICK method automatically. Previously, NICK binarization method set the k value between -0.2 and 0.2. This manual value is not adaptable to the document image condition. This method set k value based on a standard deviation value of an image. This method produced better binarization result than the NICK method.

G22: The method combines local and global adaptive binarization which is motivated by [18] and [19]. It extracts the ink pixels from the background using an inpainting-based background estimator and perform an image normalization. This inpainting method is based on a hybrid sparse representation. A dilated version of Niblack [20] output is employed as an inpainting mask. This method performs the global binarization on the normalized image similar to [18]. This method executes the local binarization by an adaptive image contrast comprised of ink-edge detection and local threshold estimation as like [19]. Finally, it merges the local and global binarized output and remove very small components.

E. Results and Analysis

The evaluation results for all submitted methods are presented in Table II. We calculated the average score of each evaluation measure from the 103 images of testing subset. The method proposed by G17 outperforms all other methods in FM and PSNR measures, but the method of G12 achieves the best (lowest) NRM measure. We also calculated the average score per collection and we compared the result with the best method (G2) from ICFHR 2016 competition [3] which was tested on the same Balinese palm leaf manuscript test set (see Table III). For each collection, the method proposed by G17 outperforms all other methods in most of the evaluation measures. But, the method of G2-2016 still achieves better performance for Balinese manuscripts.

TABLE IV
PALM LEAF MANUSCRIPT DATASETS FOR TEXT LINE SEGMENTATION TASK

Manuscripts	Train	Test	Dataset
Balinese	47 pages	49 pages	Extracted from AMADI_LontarSet [9]
Khmer	50 pages	200 pages	Extracted from SleukRith Set [10]
Sundanese	31 pages	30 pages	Extracted from Sunda Dataset [11]

IV. CHALLENGE B: TEXT LINE SEGMENTATION

A. Description of the Challenge

Text line segmentation is a crucial pre-processing step of most DIA pipelines. The task aims at extracting and separating text region into individual lines. Most of line segmentation approaches in the literature require that the input image is binarized. However, due to degradation, noise are often found on historical documents such as palm leaf manuscripts, the binarization task is not able to produce good enough results. The line segmentation methods that are independent of binarization task work usually directly on color/greyscale images. Therefore, the goal of this challenge is to develop a binarization-free method to segment and extract text line regions from color/greyscale images of ancient palm leaf manuscripts.

B. Dataset

The palm leaf manuscript datasets for text line segmentation task are presented in Table IV. The text line segmentation ground truth data for Balinese and Sundanese manuscripts have been generated by hand based on the binarized ground truth images [9]. For Khmer, an ID of the line which it belongs to, is associated to each annotated character. The region of a text line is the union of the areas of the polygon boundaries of all annotated characters composing it [5], [10]. The ground truth for line segmentation of the training set is provided. The ground truth of each file is a raw image file. Each pixel stores a positive integer value corresponding to the ID of the text line it belongs to. For the background or undefined region, its pixel stores a zero value. This challenge has only one single track, text line segmentation for a mixed collection of palm leaf manuscripts from Bali, Cambodia and Sunda.

C. Performance Evaluation

We use the evaluation criteria and tool provided by ICDAR2013 Handwriting Segmentation Contest³ [21]. First, the one-to-one (*o2o*) match score is computed for a region pair based on the evaluators acceptance threshold. To calculate the *o2o* match score, the MatchScore table is first calculated according to the intersection of the pixel sets of the result and the ground truth⁴. A pair of ground truth and result region is a one-to-one match (*o2o*) if the matching score for this pair is equal to or above the evaluator's acceptance threshold [22]. In our experiments, we used 90% as the acceptance threshold. Let N be the count of ground truth elements, and M be the

³<http://users.iit.demokritos.gr/~nstam/ICDAR2013HandSegmCont/Protocol.html>

⁴<http://users.iit.demokritos.gr/~nstam/ICDAR2013HandSegmCont/Evaluation.html>

TABLE V
EVALUATION RESULTS OF CHALLENGE B

Group	N	M	o2o	DR	RA	FM
G17	1,271	1,577	962	75.68%	61.00%	67.55%
Evaluation Results of Challenge B per collection						
G17-Bali	182	264	109	59.89%	41.28%	48.87%
APP-Bali [24]	182	191	164	90.10%	85.86%	87.93%
G17-Khmer	971	1153	778	80.12%	67.47%	73.25%
APP-Khmer [24]	971	990	910	93.71%	91.91%	92.80%
G17-Sunda	118	160	75	63.55%	46.87%	53.95%
APP-Sunda	-	-	-	-	-	-

count of result elements. With the *o2o* score, three metrics are calculated as described in [22], [23]: detection rate (*DR*), recognition accuracy (*RA*), and performance metric (*FM*).

D. Participant Methods

G17: The processed image from binarization challenge (see G17 in sub section III-D) is further segmented into line using modified horizontal projection profile. The projection profile is smoothened using a low pass filter to avoid spurious windows in the segmentation of lines. Windows are formed based on adaptive threshold. The components falling within the window are grouped as a single line segment from the binarized image. Because this method localizes only text areas from binarized images, to satisfy the evaluation criteria, the results are modified by recursively applying a dilation operation using a 3 by 3 cross-shaped structuring element until all background pixels are filled. The new results are obtained from the intersection between the binarization ground truth and the modified results.

E. Results and Analysis

The evaluation results and ranks for the only submitted method is presented in Table V. We also presented the results per collection and we compared the results with the Adaptive Path Finding (APP) method [5] reported in [24] which was tested on the same Balinese and Khmer palm leaf manuscript test set. The method of G17 achieves their best result on Khmer manuscripts. But in overall evaluation, the APP method still shows better performances.

V. CHALLENGE C: ISOLATED CHARACTER/GLYPH RECOGNITION

A. Description of the Challenge

In the case of segmentation based text recognition method, the isolated character recognition task is a very important process [25]. A proper feature extraction and a correct classifier selection can increase the recognition rate [26]. Although many methods for isolated character recognition have been widely developed and tested especially for Latin based scripts and alphabets, there is still a need for in-depth evaluation of those methods to be applied for various other types of scripts with optimal performance. It includes the isolated character recognition task for many Southeast Asian scripts, and more specifically the scripts which were written on the ancient palm leaf manuscripts.

TABLE VI
PALM LEAF MANUSCRIPT DATASETS FOR ISOLATED CHARACTER/GLYPH RECOGNITION TASK

Manuscripts	Classes	Train (images)	Test images)	Dataset
Balinese	133	11,710	7,673	AMADL_LontarSet [3], [6], [9]
Khmer	111	113,206	90,669	SleukRith Set [10]
Sundanese	60	4,555	2,816	Sunda Dataset [11]

B. Dataset

The palm leaf manuscript datasets for isolated character/glyph recognition task are presented in Table VI. The glyph annotated ground truth images creation for the manuscripts can be found in [9], [11], [13]. This challenge has three different tracks, the isolated character recognition for palm leaf manuscripts from Bali (Track 1), from Cambodia (Track 2), and from Sunda (Track 3).

C. Performance Evaluation

Following the evaluation method from ICFHR competition [3], the recognition rate, i.e., the percentage of correctly classified samples over the test samples: C/N is calculated, where C is the number of correctly recognized samples, and N is the total number of test samples.

D. Participant Methods

G20: The system is designed to handle two key problems, i.e., limited available data and large similarity among characters. First, a practical data augmentation strategy is used to solve the problem of limited available data, which prevents the model from falling into over-fitting. Second, a very deep convolutional neural network (100 layers) with dense connection is trained to classify these similar characters. It should be noted all training data come from the competition sponsor. No external data is used in this system.

G21: No description. The result was submitted only for Track 3.

E. Results and Analysis

The evaluation results for all submitted methods are presented in Table VII. We also compared the results with the best method (G1) from ICFHR 2016 competition [3] which was tested on the same Balinese palm leaf manuscript test set, and also with CNN and handcrafted feature (HoG-NPW-Kirsch-Zoning with Unsupervised Feature Learning + Neural Network) method reported in [24]. The method of G20 outperforms all other methods for each track of manuscript collections.

VI. CHALLENGE D: WORD TRANSLITERATION

A. Description of the Challenge

In order to make the palm leaf manuscripts more accessible, readable and understandable to a wider audience, an optical character recognition (OCR) system should be developed. In

TABLE VII
EVALUATION RESULTS OF CHALLENGE C (% RECOG. RATE)

Group	Track 1 Balinese	Track 2 Khmer	Track 3 Sundanese
G20	92.17	97.21	86.54
G21	No Result	No Result	84.98
G1-2016 [3]	88.39	-	-
CNN [24]	85.39	93.96	79.05
HF [24]	85.63	92.44	79.33

TABLE VIII
PALM LEAF MANUSCRIPT DATASETS FOR WORD TRANSLITERATION TASK

Manuscripts	Train (images)	Test (images)	Text	Dataset
Balinese	15,022 from 130 pages	10,475 from 100 pages	Latin	AMADL_LontarSet [3], [9]
Khmer	16,333 (part of 657 pages)	7,791 (part of 657 pages)	Latin and Khmer	SleukRith Set [10]
Sundanese	1,427 from 20 pages	318 from 10 pages	Latin	Sunda Dataset [11]

many DIA systems, word or text recognition is the final task on the processing pipeline. But, normally in Southeast Asian script the speech sound of the syllable changes according to some certain phonological rules. In this case, an OCR system is not sufficient. Therefore, a transliteration system should also be developed to help to transliterate the ancient scripts on these manuscripts. By definition, transliteration is defined as the process of obtaining the phonetic translation of names across languages [27]. Transliteration involves rendering a language from one writing system to another. In [27], the problem is stated formally as a sequence labeling problem from one language alphabet to other. It helps to index and to access quickly and efficiently to the content of the manuscripts. In this challenge, the methods are evaluated to transliterate the words from three different scripts of palm leaf manuscript.

B. Dataset

The palm leaf manuscript datasets for word transliteration task are presented in Table VIII. For Khmer dataset, all characters in the page have been annotated and grouped together into words. More than one label may be given to the created word. The order of how each character in the word is selected is also kept [10]. Balinese and Sundanese word dataset were manually annotated by using Aletheia [28]. This challenge has four different tracks, the word transliteration for palm leaf manuscripts from Bali (Track 1), from Cambodia (Track 2), from Sunda (Track 3), and from a mixed collection of palm leaf manuscripts from Bali, Cambodia and Sunda (Track 4).

C. Performance Evaluation

The Character Error Rate (CER) is calculated. The CER is defined by edit distance metric between ground truth and recognizer output and is computed by using the provided OCRopy function `ocropus-errs`⁵. This distance is defined as the

⁵<https://github.com/tmbdev/ocropy/blob/master/ocropus-errs>

ration of insertion, deletion and substitution actions compared to the total length of the word [29].

D. Participant Methods

G3: This method uses a CNN-RNN hybrid type network for the transcription process. Raw grayscale pixel values are used as input to the network. This network consists of three components: A spatial transformer layer, a Resnet-18 like convolutional feature extractor and finally two BLSTM layers. Connectionist Temporal Classification (CTC) loss function is used to avoid segmenting the word image and to predict the sequence of ascii characters present in the given image. Vocabulary from Leipzig corpora is used to synthesize synthetic training data. First, a model which combines all the training data present in the three tracks is trained. After its convergence, synthetic data for the particular track (1-3) are used and pre-train the model to that. After that, the model is finally fine-tuned to the train data of respective track. For the 4th track, a script detection model using the same architecture as above is also trained, which classifies an image into one of the three languages. As per the output of the script detection model, the models from the other three tracks are used. Various data-augmentation techniques such as using a combination of elastic distortion, various affine transforms (shearing, padding, etc.) and multi-scale image augmentation are also used.

G13: As the word transliteration for southeast Asian palm leaf manuscripts can be seen as an image to sequence task, the encoder-decoder approach is used. A convolutional neural network encoder that takes palm leaf manuscripts as image input is employed. A densely connected convolutional network is adopted as it can strengthen feature extraction and facilitate gradient propagation especially on a small training set. A recurrent neural network decoder equipped with an attention mechanism is employed to generate the transliteration sequence. The recurrent neural network is based on gated recurrent units to alleviate gradient vanishing and exploding problem. The correspondence between the input manuscripts and the output transliteration sequences is learned automatically by the attention mechanism.

E. Results and Analysis

The evaluation results and ranks for all submitted methods are presented in Table IX. We also compared the result with the LSTM method reported in [24] which was tested on the same Balinese palm leaf manuscript test set. The method of G13 outperforms all other methods for each track of manuscript collections.

VII. CONCLUSIONS

In this Competition on Document Image Analysis Tasks for Southeast Asian Palm Leaf Manuscripts, the winner for Challenge A (Binarization) is Deepak Kumar, from Department of Electronics & Communication Engineering, Dayananda Sagar Academy of Technology and Management (DSATM),

TABLE IX
EVALUATION RESULTS OF CHALLENGE D (% CER)

Group	Track 1 Balinese	Track 2 Khmer	Track 3 Sundanese	Track 4 Mixed
G3	11.59	4.51	9.68	7.16
G13	9.54	3.38	8.81	5.62
LSTM [24]	39.70	-	-	-

Bengaluru, India. The winner for Challenge C (Isolated character/glyph recognition) is Zi-Rui Wang, Jun Du and Wen-Chao Wang, from University of Science and Technology of China, China. The winner for Challenge D (Word transliteration) is Jianshu Zhang, Jun Du, and Lirong Dai, from National Engineering Laboratory for Speech and Language Information Processing, University of Science and Technology of China, China. For challenge with only one participant, we have decided that the participant is a winner if the result of the proposed method is better than the baseline method. In the case of the challenge B, it was not the case. The dataset of each challenge is available on http://amadi.univ-lr.fr/ICFHR2018_Contest/.

ACKNOWLEDGMENT

The authors would like to thank Museum Gedong Kertya, Museum Bali, Mr. Undang Ahmad Darsa, the philologists from Sundanese Centre Studies of UNPAD, the Situs Kabuyutan Ciburuy Garut, all families in Bali, Indonesia, the EFEO team, the Buddhist Institute, and the National Library in Cambodia for providing us with samples of palm leaf manuscripts. We also thank the students from the Department of PTI and the Department of Balinese Literature, UNDIKSHA, the ITC Cambodia, and the National Institute of Post, Telecommunication and ICT. This work is supported by the DIKTI BPPLN, the STIC Asia Program implemented by the French Ministry of MAEDI, and ARES-CCD (program AI 2014-2019) under the funding of Belgian university cooperation, and DRPMI UNPAD, DIKTI International Collaboration and Publication grant 2017.

REFERENCES

- [1] M W A Kesiman, S Prum, J-C Burie, and J-M Ogier. An Initial Study On The Construction Of Ground Truth Binarized Images Of Ancient Palm Leaf Manuscripts. In *13th ICDAR*, Nancy, France, August 2015.
- [2] M W A Kesiman, S Prum, I M G Sunarya, J-C Burie, and J-M Ogier. An Analysis of Ground Truth Binarized Image Variability of Palm Leaf Manuscripts. In *5th IPTA 2015*, pages 229–233, Orleans, France, November 2015.
- [3] J-C Burie, M Coustaty, S Hadi, M W A Kesiman, J-M Ogier, E Paulus, K Sok, I M G Sunarya, and D Valy. ICFHR 2016 Competition on the Analysis of Handwritten Text in Images of Balinese Palm Leaf Manuscripts. In *15th ICFHR 2016*, pages 596–601, Shenzhen, China, October 2016.
- [4] M W A Kesiman, D Valy, J-C Burie, E Paulus, I M G Sunarya, S Hadi, K Sok, and J-M Ogier. Southeast Asian palm leaf manuscript images: a review of handwritten text line segmentation methods and new challenges. *Journal of Electronic Imaging*, 26(1):011011, November 2016.
- [5] D Valy, M Verleysen, and K Sok. Line Segmentation for Grayscale Text Images of Khmer Palm Leaf Manuscripts. In *7th IPTA 2017*, Montreal, Canada, December 2017.

- [6] M W A Kesiman, S Prum, J-C Burie, and J-M Ogier. Study on Feature Extraction Methods for Character Recognition of Balinese Script on Palm Leaf Manuscript Images. In *23rd ICPR*, Cancun, Mexico, December 2016.
- [7] M W A Kesiman, J-C Burie, and J-M Ogier. A Complete Scheme Of Spatially Categorized Glyph Recognition For The Transliteration Of Balinese Palm Leaf Manuscripts. In *14th ICDAR*, Kyoto, Japan, November 2017.
- [8] Byron Leite Dantas Bezerra, editor. *Handwriting: recognition, development and analysis*. Computer science, technology and applications. Nova Science Publishers, Inc, Hauppauge, New York, 2017.
- [9] M W A Kesiman, J-C Burie, J-M Ogier, G N M A Wibawantara, and I M G Sunarya. AMADI_lontaset: The First Handwritten Balinese Palm Leaf Manuscripts Dataset. In *15th ICFHR 2016*, pages 168–172, Shenzhen, China, October 2016.
- [10] D Valy, M Verleysen, S Chhun, and J-C Burie. A New Khmer Palm Leaf Manuscript Dataset for Document Analysis and Recognition SleukRith Set. In *4th International Workshop on HIP*, Kyoto, Japan, November 2017.
- [11] M Suryani, E Paulus, S Hadi, U A Darsa, and J-C Burie. The Handwritten Sundanese Palm Leaf Manuscript Dataset From 15th Century. In *14th ICDAR*, Kyoto, Japan, November 2017.
- [12] M. Shehzad Hanif Naveed Bin Rais. Adaptive thresholding technique for document image analysis, pages 61–66. IEEE, 2004.
- [13] D Valy, M Verleysen, and K Sok. Line Segmentation Approach for Ancient Palm Leaf Manuscripts using Competitive Learning Algorithm. In *15th ICFHR 2016*, Shenzhen, China, October 2016.
- [14] B. Gatos, K. Ntirogiannis, and I. Pratikakis. DIBCO 2009: document image binarization contest. *ICDAR*, 14(1):35–44, March 2011.
- [15] B Epshtein, E Ofek, and Y Wexler. Detecting text in natural scenes with stroke width transform. In *2010 IEEE CVPR*, pages 2963–2970. IEEE, June 2010.
- [16] Nicholas R. Howe. A Laplacian Energy for Document Binarization. In *11th ICDAR*, pages 6–10. IEEE, September 2011.
- [17] K Saddami, K Munadi, S Muchallil, and F Arnia. Improved Thresholding Method for Enhancing Jawi Binarization Performance. In *14th ICDAR*, pages 1108–1113. IEEE, November 2017.
- [18] K. Ntirogiannis, B. Gatos, and I. Pratikakis. A combined approach for the binarization of handwritten document images. *Pattern Recognition Letters*, 35:3–15, January 2014.
- [19] Bolan Su, Shijian Lu, and Chew Lim Tan. Robust Document Image Binarization Technique for Degraded Document Images. *IEEE Transactions on Image Processing*, 22(4):1408–1417, April 2013.
- [20] Wayne Niblack. *An introduction to digital image processing*. Prentice-Hall International, Englewood Cliffs, N.J, 1986.
- [21] N Stamatopoulos, B Gatos, G Louloudis, U Pal, and A Alaci. ICDAR 2013 Handwriting Segmentation Contest. pages 1402–1406. IEEE, August 2013.
- [22] G. Louloudis, B. Gatos, I. Pratikakis, and C. Halatsis. Text line and word segmentation of handwritten documents. *Pattern Recognition*, 42(12):3169–3183, December 2009.
- [23] I.T. Phillips and A.K. Chhabra. Empirical performance evaluation of graphics recognition systems. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 21(9):849–870, September 1999.
- [24] M W A Kesiman, D Valy, J-C Burie, E Paulus, M Suryani, S Hadi, M Verleysen, S Chhun, and J-M Ogier. Benchmarking of Document Image Analysis Tasks for Palm Leaf Manuscripts from Southeast Asia. *Journal of Imaging*, 4(2):43, February 2018.
- [25] A Aggarwal, K Singh, and K Singh. Use of Gradient Technique for Extracting Features from Handwritten Gurmukhi Characters and Numerals. *Procedia Computer Science*, 46:1716–1723, 2015.
- [26] M Z Hossain, M A Amin, and H Yan. Rapid Feature Extraction for Optical Character Recognition. *CoRR*, abs/1206.0238, 2012.
- [27] P Shishtla, V S Ganesh, S Subramaniam, and V Varma. A language-independent transliteration schema using character aligned models at NEWS 2009. page 40. Association for Computational Linguistics, 2009.
- [28] C. Clausner, S. Pletschacher, and A. Antonacopoulos. Aletheia - An Advanced Document Layout and Text Ground-Truthing System for Production Environments. In *2011 ICDAR*, pages 48–52. IEEE, September 2011.
- [29] M Jenckel, S S Bukhari, and A Dengel. anyOCR: A sequence learning based OCR system for unlabeled historical documents. In *23rd ICPR*, pages 4035–4040. IEEE, December 2016.