

Towards Letter Shape Prior and Paleographic Tables Estimation in Hebrew First Temple Period Ostraca

Arie Shaus

Department of Applied Mathematics
Tel Aviv University, Tel Aviv, Israel
ashaus@post.tau.ac.il

Eli Turkel

Department of Applied Mathematics
Tel Aviv University, Tel Aviv, Israel
turkel@post.tau.ac.il

ABSTRACT

The problem of finding a prototype for typewritten or handwritten characters belongs to a family of “shape prior” estimation problems. In epigraphic research, such priors are derived manually, and constitute the building blocks of “paleographic tables”. Suggestions for automatic solutions to the estimation problem are rare in both the Computer Vision and the OCR/Handwriting Text Recognition communities. We review some of the existing approaches, and propose a new robust scheme, suitable for the challenges of degraded historical documents. This fast and easy to implement method is employed for ancient Hebrew inscriptions dated to the First Temple period.

CCS CONCEPTS

• **Computing methodologies** → **Artificial intelligence** → **Computer vision** → **Computer vision problems** → **Shape inference**

KEYWORDS

letter shape prior, character templates, document-specific alphabet, glyph extraction, ideal/Platonic prototypes, allograph, epigraphy, paleographic tables, historical documents, Hebrew ostraca, First Temple period

1 INTRODUCTION

The issue of prototype inference for typewritten or handwritten characters belongs to a broad type of “shape prior” determination problems, which has gathered substantial research interest during the last two decades. Nevertheless, research deriving shape prior of handwritten or printed characters are relatively rare in both the Computer Vision (CV) and the OCR/Handwriting Text Recognition (HTR) communities. The lack

of interest of CV scientists can be explained by the specificity of this challenging problem. On the other hand, most of the HTR studies focus on producing ever improving recognition engines – a related, yet not directly dependent problem. The relatively low interest in the subject resulted in diverse terms used by the existing publications. Among the related terms are “letter/handwriting prototypes”, “document-specific alphabet”, “reconstructed font”, “glyph extraction”, “character template estimation”, “character models”, “codebook generation”, “ideal/Platonic prototypes” and “letter shape priors”. In what follows, we shall use the last term, common in the CV community.

The reconstructed priors can be utilized for issues such as denoising automatic damage removal, compression, archiving, as well as handwriting and style analyses. Moreover, in the context of historical texts, the priors are closely related to the so called “paleographic tables” – a basic and crucial instrument in the toolbox of the historical epigrapher (an expert on ancient writings). Commonly, such tables contain one characteristic example of each letter type for each inscription in a given corpus; see example on Fig. 1. The tables are used to trace the similarities and the differences within the handwriting of different localities and time periods. This labor-intensive process joins other manually performed epigraphic tasks. Indeed, currently, the imaging, the creation of the facsimile (a black and white depiction of the inscription), the recognition of the letters, the transcription, the creation of paleographic tables, as well as their analysis are all carried out manually by epigraphic experts. Such an effort is extremely time-consuming, producing results which may accidentally mix-up documentation with interpretation. In other words, the quality of the paleographic tables is often debated and, unfortunately, cannot be treated as an established “ground truth”.

In previous publications, we dealt with imaging techniques of ancient ostraca (ink on clay inscriptions; mostly dated to the 7th century BCE) [1,2], as well as their binarizations [3-6] and writers’ identifications [7,8]. The current research is a continuation of these studies. Its envisioned objective is an automatically derived paleographic table, accompanied by its algorithmic analysis. In this paper, we will concentrate on a challenging intermediate goal of obtaining the main building block of such a table, i.e. the letter shape prior.

For consistency purposes, the following terminology is used throughout this article. By “letters” we designate the members of the alphabet, e.g. “aleph”, “bet”, etc. Their realizations by the writer are the particular characters, e.g. an inscription may

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from Permissions@acm.org.

HIP2017, November 10–11, 2017, Kyoto, Japan

© 2017 Copyright is held by the owner/author(s). Publication rights licensed to ACM.

ACM ISBN 978-1-4503-5390-8/17/11 \$15.00

<https://doi.org/10.1145/3151509.3151511>

contain several “bet” characters. A “letter shape prior”, or in short “letter prior”, represents a typical way of depicting a given letter.

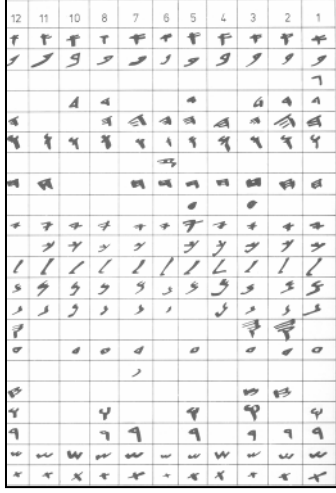


Figure 1: Manually created paleographic table, recording “typical” representatives for each letter in the alphabet (adapted from [25]).

2 PRIOR ART

Ostensibly, the task of estimating the letter prior seems to be relatively straightforward, requiring a registration of the character images, their accumulation and subsequent thresholding. However, in reality, this undertaking turns out to be surprisingly difficult. Indeed, elastic image registration is an NP-complete problem [9]. Moreover, multiple template alignment estimation was also shown to be NP-complete [10]. Thus, the existing solutions of mutual registration problem are heuristic, and tend to balance between the computational costs and the quality of the result.

Kopec and Lomelin [10] proposed a sophisticated Aligned Template Estimation (ATE) framework, in which overlapping glyphs templates were searched in a page image. The authors used a two-phase iterative training algorithm, encompassing an alignment of pre-existing transcriptions given an initial guess (existing transcriptions), as well as an ATE stage. The ATE step was implemented via a likelihood maximization procedure. The technique was designed for typewritten characters. Its results were reasonable given sufficiently large data and a number of iterations. Nevertheless, some artifacts were present in the resulting “priors”, due to the method’s “unawareness” of the different character properties, and inexact segmentation boundaries. Bern and Goldberg [11] proposed a variation on the theme of super-resolution within a single image, also in the context of printed text. Given a relatively clean binarized document image, the letters were registered, then iteratively clustered, taking phenomena such as touching letters into account. A Bayesian calculation yielded a prior, which was utilized for image de-noising purposes. The results of this algorithm also exhibited certain artifacts, due to the exceedingly

fine-grained clustering, and mistaking noisy characters for distinct glyphs.

For handwriting, several papers included prior estimation as an intermediate step in handwriting synthesis (i.e. a simulation of a particular handwriting style given a few writing examples). As opposed to the relatively fixed typewritten characters of previous works, now a more challenging cursive writing, with its high variance, was considered. The inputs in these cases were clean and thinned writing examples. In [12], after a segmentation achieved by a Hidden Markov Model, a curve control point interpolation was performed. Wang et al. [13] extracted priors in addition to a “tri-unit” technique (akin to the tri-grams of Speech Recognition). This was used in order to identify different types of “contact” strokes between various characters. The shape prior creation was composed of control point extraction (Gabor filters leading to a B-splines approximation), affine registration and shape prior parameter estimation stages, with impressive results.

Edwards and Forsyth [14] derived a shape prior in the complicated world of historical documents (12th century manuscript). The authors initiated the priors with hand cut examples. The page image was then segmented into words and characters; each word possessed several possible segmentations (represented by a graph). For each word, the different possible segmentations were searched within a pre-existing dictionary (in the target language) by comparing the word image with the candidate word image derived from shape priors. The high confidence matches were accepted, and then the shape priors were updated. If necessary, new shape priors (possibly more than one for a single character) were created. The process was then repeated. A similar statistical language model was also utilized in [15,16], where candidate words are checked vs. an English corpus. Words (token) co-occurrence statistics was used in order to correctly identify problematic characters.

A noteworthy modern variational approach in a historic setting was presented in [17,18]. Given a set of character edges, a confidence map (shape prior) was created for each character individually via a Gradient Vector Flow. Subsequently, the confidence map could be fitted back into the document image, utilizing the Active Contour method, in order to achieve high-quality segmentation. Panagopoulos et al. [19] utilized estimated “ideal” or “Platonic” prototypes for each letter of historical inscriptions for the purpose of writer identification analysis.

3 THE WRITING MEDIUM AND THE PROPOSED ALGORITHM

This paper deals with ancient Hebrew ostraca (ink on clay inscriptions), created at the end of the First Temple period, ca. 600 BCE. These texts, written in alphabetical Paleo-Hebrew writing, are of mundane nature, covering issues such as food supplies and movement of troops. Many of the ostraca were not composed by professional scribes [7], and therefore the variability of the handwriting is very high. The inscriptions are quite short (typically containing 30-100 characters), and their state of preservation is poor (the ostraca are often broken, and parts of the writing are soiled).

These characteristics of the writing medium influenced the design of our algorithm. Contrary to prior art, only small amounts of characters for each type of letter are present for each ostrakon. Moreover, the inscriptions are highly degraded (with blurred and erased characters, as well as cracks and stains easily mistaken for characters). Hence, we preferred robust statistical estimators such as median and medoid (a representative object, whose dissimilarity to other objects in the population is minimal) over the commonly used mean, which is easily susceptible to noise (for another use of medoids and medians in a related setting, see [4]).

We assume grayscale images of the ostraca (e.g. acquired by methods described in [1,2]). We also pre-suppose imperfect black and white facsimiles, registered to the grayscale ostraca images. Such facsimiles are often created by epigraphers (an automatic creation of facsimiles, i.e. binarization, can also be attempted, see [3,4]). The facsimiles (manual depictions of the inscription) are only utilized for preliminary segmentation purposes, in a manner similar to that described in [3], i.e. the registered facsimiles provide us with an initial indication regarding the position and the type of inscriptions' characters within the ostraca images. The algorithm utilizes the cropped (dilated and padded) character grayscale images; chooses a medoid image via simple registration procedure; registers all the other character images to the medoid image; calculates the initial prior via median calculation per each pixel coordinate; thresholds the prior via modification of Otsu's algorithm [20], and if needed, smoothes the result.

The detailed steps of our algorithm, for a given inscription and letter, are:

1. **Cropping character images:**
 - 1.1. The characters' convex hulls of width w_i and height h_i ($i = 1, \dots, K$), are found at the facsimile level.
 - 1.2. The convex hulls are dilated by $PAD \cdot \max\{w_i, h_i\}$ pixels (assuming 4-connectivity), with respect to a pre-determined parameter PAD (herein, $PAD = 0.1$).
 - 1.3. The locations of the dilated convex hulls in the facsimile image are used in order to crop rectangular images $S_i(m, n) : [1, M_i] \times [1, N_i] \rightarrow [0, 255]$ of the characters from the grayscale ostraca images. Pixels corresponding to the dilated convex hulls assume the grayscale values of the inscription image, while other pixels assume the padding value of 255.
2. **Padding character images:**
 - 2.1. The maximal dimensions of the character images are calculated: $M = \max\{M_i\}$, $N = \max\{N_i\}$.
 - 2.2. These dimensions are utilized in order to create padded character images of common size. The padding (by 255) is applied symmetrically on the opposite sides of S_i , resulting in $P_i(m, n) : [1, M] \times [1, N] \rightarrow [0, 255]$, images of the same size.
3. **Initial characters' registration:**
 - 3.1. For each $i = 1, \dots, K$ and for each $j = 1, \dots, K$ s. t. $i \neq j$ a normalized cross-correlation fit ρ_{ij} [21] is calculated between P_i and S_j .
 - 3.2. The (not necessarily symmetrical) distances d_{ij} are calculated: $d_{ij} = \sqrt{(1 - \rho_{ij}) / 2}$ (see [22] for details).
 - 3.3. A medoid index $l = \arg \min_i \left(\sum_j d_{ij} \right)$ and an initial registered image $R_l = P_l$ are established.
 - 3.4. For all $i = 1, \dots, K$, s. t. $i \neq l$, the $S_i(m, n)$ images are translated according to their optimal shift with respect to R_l (calculated at stage 3.1), in order to obtain registered images R_i ; their padding value is 255.
4. **Letter prior initialization:**
The initial prior L_{init} is calculated via median for each pixel coordinate, over all the registered character images:
$$L_{init}(m, n) = \underset{i=1, \dots, K}{\text{median}} \{R_i(m, n)\}.$$
5. **Letter prior thresholding:**
A thresholded prior image L_{thr} is calculated via $L_{thr} = \text{Otsu}^*(L_{init})$, where Otsu^* is an adaptation of Otsu's algorithm [20] ignoring the histogram value of 255 (i.e. the padding values of steps 1.3, 2.2 and 3.4, which might skew the statistics).
6. **Letter prior smoothing:**
A smoothed prior image L_{sm} is calculated via $L_{sm} = \text{MorphCV}(L_{thr}, REG)$, where MorphCV is a morphological solution to the popular Chan-Vese [23] framework, introduced and analyzed in [24]. The latter demonstrates the equivalence of variational and median-based smoothing. REG as an optional regularization (smoothing) parameter, controlling the median filter radius.
7. **Optional letter prior calculation loop:**
The estimated prior L_{sm} can now be plugged-in at step 3.4, with all the S_i optimally fitted to L_{sm} instead of the medoid P_l . The resulting collection can then be refined (via the median, as in step 4), the outcome thresholded by Otsu^* (as in step 5), and its result smoothed via MorphCV (as in step 6). The loop can be either stopped at this stage, or repeated until convergence.

4 RESULTS

Experimentations with the proposed framework were conducted on three relatively large ostraca, belonging to the First Temple period corpus of Hebrew inscription from the Arad fortress [25], dated to ca. 600 BCE. In particular, we tested different configurations of our method on Arad 1, Arad 2 and Arad 24b (verso side) inscriptions. The 8-bit grayscale images of the ostraca were approximately of the same resolution, with a typical character size of 30,000-60,000 pixels (width and height varying depending on the character). Registered facsimiles, colored according to letter types, were also utilized; see Figs. 2-4 for images of ostraca and their facsimiles.

This size of the ostracon images was reduced by half (on each side) in some of the experiments, in order to test the performance of the algorithm in such a setting. In total, 310 characters were utilized. Several representative examples of the algorithm's steps and its outcomes are provided below.

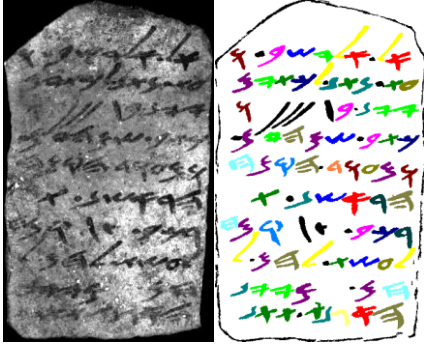


Figure 2: Arad 1 - an ostracon image and its corresponding facsimile. The various colors of the facsimile (adapted from [25]) indicate different letter types.

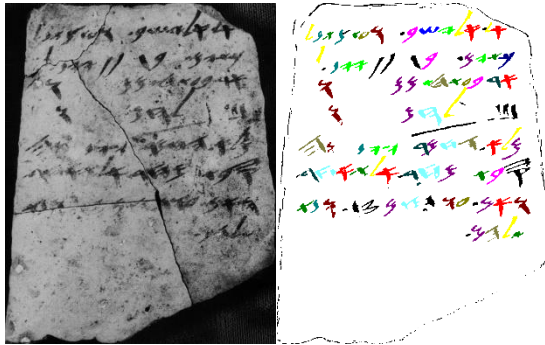


Figure 3: Arad 2 - an ostracon image and its corresponding facsimile. The various colors of the facsimile (adapted from [25]) indicate different letter types.

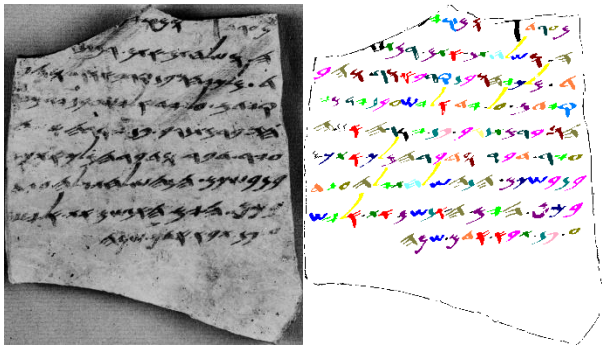


Figure 4: Arad 24b - an ostracon image and its corresponding facsimile. The various colors of the facsimile (adapted from [25]) indicate different letter types.

Fig. 5 shows an illustration of the algorithm's flow on a letter "yod" from Arad 24b. On the top row, a refinement of the prior (based on 14 characters) is shown, with no attempt at regularization (smoothing). On the bottom row, three consecutive priors are regularized by an algorithm [24], performing median-based smoothing with median filter radius set to $REG=5$. Similarly, Fig. 6 shows the steps for a regularized computation of "mem" from Arad 2 (based on 10 characters).



Figure 5: An example of the algorithm's flow for "yod" letter, Arad 24b. Top: a median-based initialization of a prior (utilizing information from 14 characters), and an estimation of two consequent priors, with no attempt at regularization (smoothing). Bottom: three consecutive priors are regularized with $REG=5$.



Figure 6: Steps for a regularized prior computation of "mem" from Arad 2 (based on 10 characters).

Fig. 7 provides a computation of a prior for the letter "ayin" from Arad 1 ostracon, in both full and partial resolution (subsequently scaled to the same size). It can be observed that in this case, "less is more", with higher resolution input resulting in unwarranted artifacts, mistaken for delicate features.



Figure 7: The letter "ayin" from Arad 1 (based on 3 characters). Top: computation of letter prior for full resolution imagery, $REG=5$. Bottom: computation of letter prior for partial resolution (halved in each axis), with no regularization, $REG=5$ and $REG=10$.

As visual observations of the results are subjective in nature, and since neither ancient nor modern writing specimens possess

a reliable and uncontested ground truth for *letters' priors* (in fact, even the facsimiles utilized herein tend to be rather imprecise [26]), we settled on an experimental methodology akin to the one presented in [6]. Every *facsimile character* of every *ostracon* was treated (in its turn) as “artificial” ground truth for a *letter's prior*. Subsequently, “synthetic” character instances were obtained by adding incrementally increasing levels of disturbances to this image, resulting in different grayscale images. These were utilized to infer a prior. Finally, this estimation was compared to the “ground truth”, in order to deduce the precision and recall. Some details on the settings of various experiments are provided in Table 1.

Table 1: Experiments' Settings

Experi- ment	Settings		
	<i>Gaussian noise levels</i>	<i>Number of instances for each prior</i>	<i>Total number of experiments</i>
#1	Standard deviation of 200 gray values (out of 255).	2, 4, 6, 8, 10	1550
#2	Standard deviation of 50, 100, 150, 200 and 250 gray values (out of 255).	5	1550

In total, **3100** experiments were conducted. The whole series of experiments took 586.2 seconds on an Intel Core M-5Y10c 0.8Ghz, with 8 GB of memory on a single thread with no parallel computing.

The results of experiment #1 for different ostraca can be seen in Tables 2-4. They indicate the robustness of the algorithm with respect to the number of characters, with good results for at least 4 characters.

The results of experiment #2 for different ostraca can be seen in Tables 5-7. They indicate only a minor influence of the amount of noise on the average precision and recall.

Table 2: Results of Experiment #1 for Arad 1 Ostracon

Gaussian noise levels	Results for each scenario		
	<i>Number of character instances for each prior</i>	<i>Average precision</i>	<i>Average recall</i>
std = 200 gray values (out of 255).	2	94.55%	88.13%
	4	98.55%	98.17%
	6	99.07%	98.88%
	8	99.18%	99.02%
	10	99.19%	99.07%

Table 3: Results of Experiment #1 for Arad 2 Ostracon

Gaussian noise levels	Results for each scenario		
	<i>Number of character instances for each prior</i>	<i>Average precision</i>	<i>Average recall</i>
std = 200 gray values (out of 255).	2	90.28%	89.00%
	4	97.64%	97.71%
	6	98.46%	98.39%
	8	98.63%	98.50%
	10	98.66%	98.52%

Table 4: Results of Experiment #1 for Arad 24b Ostracon

Gaussian noise levels	Results for each scenario		
	<i>Number of character instances for each prior</i>	<i>Average precision</i>	<i>Average recall</i>
std = 200 gray values (out of 255).	2	87.23%	89.34%
	4	97.44%	97.82%
	6	98.73%	98.64%
	8	99.04%	98.87%
	10	99.14%	98.96%

Table 5: Results of Experiment #2 for Arad 1 Ostracon

Number of instances	Results for each scenario		
	<i>Gaussian noise level (std)</i>	<i>Average precision</i>	<i>Average recall</i>
5	50	99.05%	99.03%
	100	99.11%	99.05%
	150	99.27%	98.92%
	200	99.01%	98.10%
	250	97.83%	95.98%

Table 6: Results of Experiment #2 for Arad 2 Ostracon

Number of instances	Results for each scenario		
	<i>Gaussian noise level (std)</i>	<i>Average precision</i>	<i>Average recall</i>
5	50	98.43%	98.42%
	100	98.53%	98.45%
	150	98.76%	98.35%
	200	98.45%	97.52%
	250	96.81%	95.43%

Table 7: Results of Experiment #2 for Arad 24b Ostrakon

Number of instances	Results for each scenario		
	Gaussian noise level (std)	Average precision	Average recall
5	50	99.09%	99.05%
	100	99.14%	99.05%
	150	99.19%	98.71%
	200	98.62%	97.56%
	250	96.81%	95.27%

5 CONCLUSIONS AND FUTURE DIRECTIONS

The results of the experiments indicate the potential of our technique, particularly in the context of degraded historical characters. The algorithm is straightforward to implement, and is very fast. The dependence of our method on the number of characters is limited, and the results are only moderately affected by the accumulated noise.

The outcomes of the algorithm may benefit from more aggressive input filtering (e.g. by methods such as [5,6,27]). Further enhancements, worth considering, include an introduction of “weights” into the refinement process and a multiplicity of priors for a single letter in case of a high variance within the writing. The experimental section may benefit from adding further noise models.

ACKNOWLEDGMENTS

The research received initial funding from the European Research Council under the European Community's Seventh Framework Programme (FP7/2007-2013)/ERC grant agreement no. 229418, and by an Early Israel grant (New Horizons project), Tel Aviv University. This study was also supported by a generous donation from Mr. Jacques Chahine, made through the French Friends of Tel Aviv University. Arie Shaus is grateful to the Azrieli Foundation for the award of an Azrieli Fellowship. The kind assistance of Dr. Shirly Ben-Dor Evian, Ms. Sivan Einhorn, Ms. Shira Faigenbaum-Golovin, Dr. Anat Mendel Geberovich, and Mr. Barak Sober is greatly appreciated. Ostrakon images: courtesy of the Institute of Archaeology, Tel Aviv University; and of the Israel Antiquities Authority. Paleographic tables and facsimiles are courtesy of the Israel Exploration Society [25].

REFERENCES

- [1] S. Faigenbaum, B. Sober, A. Shaus, M. Moinester, E. Piasetzky, G. Bearman, M. Cordonsky, and I. Finkelstein. 2012. Multispectral Images of Ostraca: Acquisition and Analysis. *J. of Archaeological Science* 39, 12, 3581-3590, <https://doi.org/10.1016/j.jas.2012.06.013>
- [2] S. Faigenbaum-Golovin, A. Mendel-Geberovich, A. Shaus, B. Sober, M. Cordonsky, D. Levin, M. Moinester, B. Sass, E. Turkel, E. Piasetzky, I. Finkelstein. 2017. Multispectral Imaging Reveals Biblical-Period Inscription Unnoticed for Half a Century. *PLOS ONE* 12, 6, e0178400, <https://doi.org/10.1371/journal.pone.0178400>
- [3] A. Shaus, E. Turkel, and E. Piasetzky. 2012. Binarization of First Temple Period Inscriptions: Performance of Existing Algorithms and a New Registration Based Scheme. In *Proc. 13th Int. Conf. on Frontiers in Handwriting Recognition (ICFHR 2012)*, 645-650, <https://doi.org/10.1109/ICFHR.2012.187>
- [4] A. Shaus, B. Sober, E. Turkel and E. Piasetzky. 2013. Improving Binarization via Sparse Methods. In *Proc. 16th Int. Graphonomics Society Conf. (IGS 2013)*, 163-166.
- [5] S. Faigenbaum, A. Shaus, B. Sober, E. Turkel, and E. Piasetzky. 2013. Evaluating Glyph Binarizations Based on their Properties. In *Proc. 13th ACM Symp. on Document Engineering (DocEng 2013)*, 127-130, <https://doi.org/10.1145/2494266.2494318>
- [6] A. Shaus, B. Sober, E. Turkel and E. Piasetzky. 2016. Beyond the Ground Truth: Alternative Quality Measures of Document Binarizations, In *Proc. 15th Int. Conf. on Frontiers in Handwriting Recognition (ICFHR 2016)*, 495-500, <https://doi.org/10.1109/ICFHR.2016.0097>
- [7] S. Faigenbaum-Golovin, A. Shaus, B. Sober, D. Levin, N. Na'aman, B. Sass, E. Turkel, E. Piasetzky, I. Finkelstein. 2016. Algorithmic Handwriting Analysis of Judah's Military Correspondence Sheds Light on Composition of Biblical Texts. *Proceedings of the National Academy of Sciences* 113, 17, 4664-4669, <https://doi.org/10.1073/pnas.1522200113>
- [8] A. Shaus and E. Turkel. 2017. Writer Identification in Modern and Historical Documents via Binary Pixel Patterns, Kolmogorov-Smirnov Test and Fisher's Method. *J. of Imaging Science and Technology* 61, 1, 104041-104049, <https://doi.org/10.2352/J.ImagingSci.Technol.2017.61.1.010404>
- [9] D. Keyers and W. Unger. 2003. Elastic Image Matching is NP-Complete. *Pattern Recognition Letters* 24, 1, 445-453, [https://doi.org/10.1016/S0167-8655\(02\)00268-4](https://doi.org/10.1016/S0167-8655(02)00268-4)
- [10] G. E. Kopec and M. Lomelin. 1997. Supervised Template Estimation for Document Image Decoding. *IEEE Trans. Pattern Anal. Mach. Intell.* 19, 12, 1313-1324, <https://doi.org/10.1109/34.643891>
- [11] M. Bern and D. Goldberg. 2000. Scanner-Model-Based Document Image Improvement. In *Proc. 2000 Int. Conf. on Image Processing (ICIP 2000)*, 582-585, <https://doi.org/10.1109/ICIP.2000.899497>
- [12] L. Devroye and M. McDougall. 1995. Random Fonts for the Simulation of Handwriting. *Electronic Publishing* 8, 4, 281-294.
- [13] J. Wang, C. Wu, Y. Q. Xu, H. Y. Shum, and L. Ji. 2002. Learning-Based Cursive Handwriting Synthesis. In *Proc. 8th Int. Workshop on Frontiers of Handwriting Recognition (IWFHR 2002)*, 157-162, <https://doi.org/10.1109/IWFHR.2002.1030902>
- [14] J. Edwards and D. Forsyth. 2006. Searching for Character Models. In *Advances in Neural Information Processing Systems 18 (NIPS 2005)*, 331-338.
- [15] M. P. Cutter, J. van Beusekom, F. Shafait, and T. M. Breuel. 2010. Unsupervised Font Reconstruction Based on Token Co-Occurrence. In *Proc. 10th ACM Symp. on Document Engineering (DocEng 2010)*, 143-150, <https://doi.org/10.1145/1860559.1860589>
- [16] M. P. Cutter, J. van Beusekom, F. Shafait, and T. M. Breuel. 2011. Font Group Identification Using Reconstructed Fonts. *Proc. SPIE 7874, Document Recognition and Retrieval XVIII (DRR XVIII)*, 78740N-1-8, <https://doi.org/10.1117/12.873398>
- [17] I. Bar Yosef, A. Mokeichev, K. Kedem, and I. Dinstein. 2008. Global and Local Shape Prior for Variational Segmentation of Degraded Historical Characters. In *Proc. 11th Int. Conf. on Frontiers in Handwriting Recognition (ICFHR 2008)*, 198-203.
- [18] I. Bar Yosef, A. Mokeichev, K. Kedem, and I. Dinstein. 2009. Adaptive Shape Prior for Recognition and Variational Segmentation of Degraded Historical Characters. *Pattern Recognition* 42, 3348-3354, <https://doi.org/10.1016/j.patcog.2008.10.005>
- [19] M. Panagopoulos, C. Papaioyannis, P. Rousopoulos, D. Dafi, and S. Tracy. 2009. Automatic Writer Identification of Ancient Greek Inscriptions. *IEEE Trans. Pattern Anal. Mach. Intell.* 31, 8, 1404-1414, <https://doi.org/10.1109/TPAMI.2008.201>
- [20] N. Otsu, A Threshold Selection Method from Gray-Level Histograms. 1979. *IEEE Trans. Syst. Man. Cybern.* 9, 1, 62-66, <https://doi.org/10.1109/TSMC.1979.4310076>
- [21] W. K. Pratt. 1974. Correlation Techniques of Image Registration. *IEEE Trans. on Aerospace and Electronic Systems* 3, 353-358, <https://doi.org/10.1109/TAES.1974.307828>
- [22] S. Van Dongen and A. J. Enright. 2012. “Metric distances derived from cosine similarity and Pearson and Spearman correlations,” *arXiv preprint*, arXiv:1208.3145.
- [23] T. F. Chan and L. Vese. 2001. Active Contours without Edges. *IEEE Trans. Image Process.* 10, 2, 266-277, <https://doi.org/10.1109/83.902291>
- [24] A. Shaus and E. Turkel. 2016. Chan-Vese Revisited: Relation to Otsu's Method and a Parameter-Free Non-PDE Solution via Morphological Framework. In *Proc. 12th Int. Symp. on Visual Computing (ISVC 2016)*, 203-212, https://doi.org/10.1007/978-3-319-50835-1_19
- [25] Yohanan Aharoni. 2008. *Arad Inscriptions*. Israel Exploration Society, Jerusalem.
- [26] A. Shaus, E. Turkel, and E. Piasetzky. 2012. Quality Evaluation of Facsimiles of Hebrew First Temple Period Inscriptions. In *Proc. 10th IAPR International Workshop on Document Analysis Systems (DAS 2012)*, 170-174, <https://doi.org/10.1109/DAS.2012.70>
- [27] A. Shaus, E. Turkel, S. Faigenbaum-Golovin, B. Sober, E. Turkel, and E. Piasetzky. 2017. Potential Contrast - A New Image Quality Measure. In *Proc. IS&T Int. Symp. on Electronic Imaging 2017, Image Quality and System Performance XIV Conference (IQSP XIV)*, 52-58, <https://doi.org/10.2352/ISSN.2470-1173.2017.12.IQSP-226>