

Shira Faigenbaum-Golovin ^{ID} | Department of Mathematics and Rhodes Information Initiative, Duke University, Durham, NC 27708 USA | E-mail: shira.golovin@math.duke.edu

Arie Shaus ^{ID} | Department of Genetics, Harvard Medical School, Boston, MA 02115 USA and also Department of Archaeology, Tel Aviv University, Tel Aviv 6997801, Israel | E-mail: ashaus@post.tau.ac.il

Barak Sober ^{ID} | Department of Statistics and Data Science, Digital Humanities, The Hebrew University of Jerusalem, Mount Scopus 91905, Israel | E-mail: barak.sober@mail.huji.ac.il

Computational Handwriting Analysis of Ancient Hebrew Inscriptions—A Survey

Abstract—Ancient texts are unique evidence providing a glimpse into the thoughts, day-to-day life, and culture of people of long-gone eras. Paleography, the study of writing, aims at documenting the inscriptions, transliterating the texts, reconstructing their historical context, and studying the evolution of writing itself. The digital revolution gave rise to computational paleography, introducing new tools of data acquisition, image processing, and machine learning. Herein, we will provide an introduction to the emerging field of computational paleography through the lens of ancient Hebrew inscriptions, dating from the Iron Age through the Middle Ages. The years that passed since their composition had a great effect on their preservation level, including blurs, stains, and erosions; moreover, some documents tend to fade in the years after their discovery. Therefore, it is of paramount importance to promptly document ancient inscriptions using the most suitable imaging techniques, such as visible, infra-red, or multispectral imaging. Image analysis and processing techniques, such as binarizations, letter segmentation, and letters' prior estimation are valuable in their own right or may serve as a stage for subsequent tasks. We will also discuss automatic handwriting analysis and writers' identification, which could shed light on the historical background of the inscriptions.

Introduction

Among the most informative pieces of evidence regarding human history is the written word. The task of the paleographer

is to analyze, document, and decipher ancient inscriptions, as well as to study the evolution and variation of the scripts, putting them in a coherent chronological and spatial context (see Figure 1 for an illustration of a typical paleographic study). Accordingly, the paleographic information, encompassing various languages, scripts, and epochs, is indispensable in disentangling the broader historical puzzles.

Manual paleographic research faces several essential challenges. In particular, most of the paleographic tasks (e.g., comparing numerous characters across many corpora for dating purposes, or finding analogies in other extant texts) are time-consuming. Due to the commonly performed in-depth analysis (some paleographic publications may deal with a barely visible line of text, a word, or even a character!), this is true even for a single inscription, and it is certainly a major challenge if hundreds of documents are involved. Additionally, it can be argued that it is difficult for the paleographer to stay impartial within the realm of the analysis, and hence documentation might be conflated with interpretation.

Fortunately, in recent decades, the digital revolution has begun to permeate many fields of research in the humanities. In particular, it gave rise to computational paleography, introducing new tools of data acquisition, image processing, and machine learning to the field. This provided the machinery for many transformative studies pushing the boundaries of traditional paleography for different periods and types of scripts. However, in order to keep this article within a reasonable scope, we limit our discussion to various ancient Hebrew corpora, valuable due to their archaeological, historical, and theological significance.

The Hebrew inscriptions we consider span across diverse periods, beginning with the Biblical kingdoms of Israel and Judah in the Iron Age (ending with the destruction of the First Temple by Nebuchadnezzar in 586 BCE) through the Hellenistic and Roman eras (ending with the destruction of

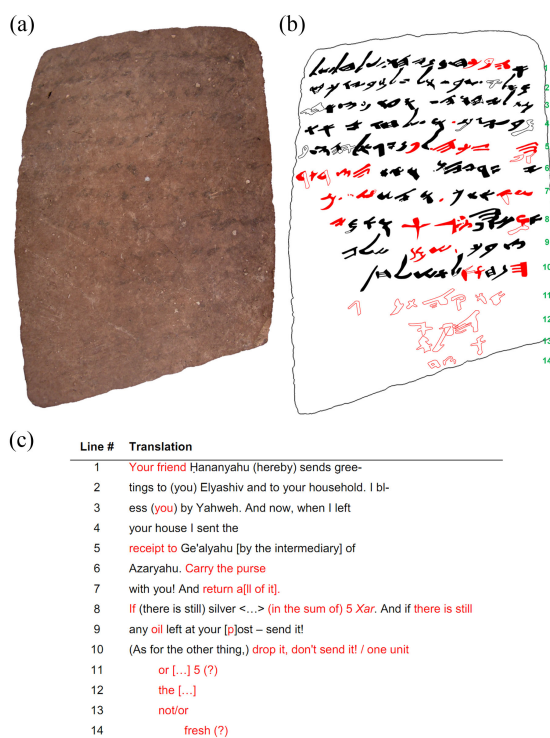


Figure 1

An example of typical paleographic tasks. (a) Ostrakon 16 (front side) from Tel Arad, ca. 600 BCE—characteristic erosions and stains are evident; (b) manual facsimile documenting the inscription; (c) translation. Adapted from [1].

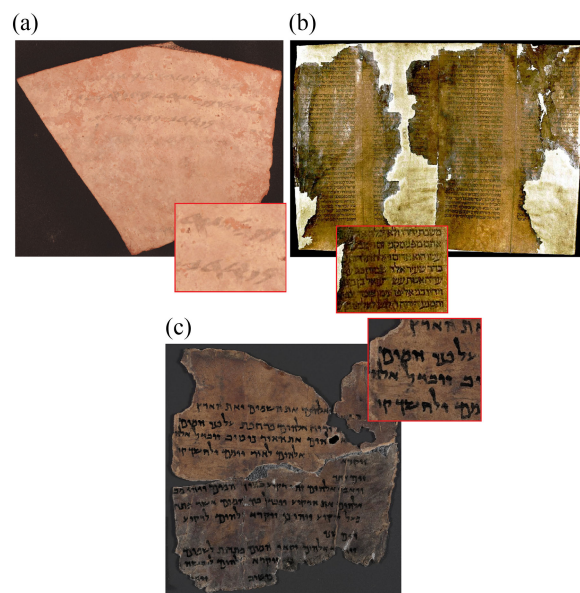


Figure 2

Ancient Hebrew inscriptions. (a) Iron Age ostraca, the Ophel ostraca, adapted from [2]. Note the blurred, eroded, and faded text. (b) Fragment of a Torah Scroll, possibly 13th century, found in the Cairo Genizah (MS. Heb. a. 4), adapted from [3]. Note the prominent stains and tears. (c) Dead Sea Scroll fragment 4Q7, Genesis 1, adapted from [4]. Note the missing parts of the inscription.

papyri, parchment, vellum, paper, and cloth. Due to the differences in time of composition and writing material, they represent various states of preservation, typically rather problematic. The inscriptions might include blurs, stains, and erosions (see Figure 2 for such issues); moreover, some documents tend to fade in the years after their discovery. The presence of such issues represents a major challenge for both classical and computational paleography.

The main impetus behind computational paleography projects is to help paleographers cope with the abundance of data, reduce the subjective involvement in the technical processes of documenting and extracting the letter shapes, and take advantage of empirical scientific methodologies to pose and resolve questions that could not have been conceived otherwise. All these tasks have to be dealt with minding the incomplete and noisy properties of the materials involved, which do not allow for the employment of off-the-shelf image processing and machine learning algorithms.

The rest of this article is organized according to common computational paleographic tasks:

- 1) *Image acquisition*: Visible, infra-red, multi- and hyper-spectral imaging, reflectance transformation imaging (RTI), Raman spectroscopy, and micro-CT.

the Second Temple by the future Roman emperor Titus in 70 CE) and ending with the prosperous and multicultural Middle Ages. The main corpora that were preserved through the years and investigated using computational methodologies are: *Hebrew ostraca* (ink-on-clay inscriptions; 8th to 6th century BCE), *the Dead-Sea Scrolls* (3rd century BCE to 1st century CE), and *the Cairo Genizah fragments* (9th to the 19th century CE); see Figure 2 for an illustration.

The Hebrew alphabet has changed significantly over the course of history. It appeared during the First Temple period, as an evolution of the proto-Canaanite alphabet (in fact, almost all modern alphabets stem from this same source). After the destruction of Jerusalem by Nebuchadnezzar II, king of Babylon, the Hebrew script has been undocumented for several centuries. When Hebrew writing re-emerges during the Second Temple period, it adopts the square Aramaic script. This script is the one used in the Dead-Sea Scrolls; its development can be seen in the Cairo Genizah fragments, as well as in modern noncursive Hebrew writing.

Additionally, the Hebrew documents under discussion herein were written on a plethora of mediums such as ostraca,

- 2) *Image analysis and processing*: Image selection, image blending, binarization, and segmentation.
- 3) *Handwriting analysis*: Writer comparison and classification, fragments' matching, letters' prior estimation, and dating based on handwriting style.
- 4) *Further tasks*: Comparison of manually created facsimiles, transcription-assistance tools, optical character recognition (OCR), and transcripts alignment.

We will conclude with our thoughts regarding the applicability of the surveyed methods for other ancient and modern types of writing, and highlight open problems and possible future research directions.

It is important to stress, that unlike other fields of research pertaining to information processing, there are no established state-of-the-art standards for either of the tasks in the pipeline. This is so because each of the research projects we discuss below is concentrated on its own set of different problems. The organization of this article is a mere attempt to put these under common titles, but in essence, the various algorithms cannot be compared as they are tailored to deal with different scenarios.

Image Acquisition

A major issue that arises in dealing with ancient inscriptions is their state of preservation. In some cases, these inscriptions were buried underground for several millennia and have gone through postdepositional processes, while in others they have been torn and traveled around the globe for centuries (e.g., Cairo Genizah texts). Exposing such delicate materials to daylight might also cause a gradual process of ink traces decay (e.g., some Iron Age ostraca which have been unearthed about 70 years ago are barely legible today). Thus, it is of utter importance to document ancient inscriptions using the most suitable imaging techniques promptly after their discovery.

Visible light photography is the most common way of documenting inscriptions. Since ancient inscriptions tend to suffer from decay, oftentimes the photographs taken shortly after their discovery are the best documentation we have. In many cases, these photographs were taken using analog cameras, and the images were recorded in the form of negatives (film and sometimes even glass). Thus, a preliminary process to enable computational analysis involves the digitization of the data (i.e., scanning the negatives). Both the Cairo Genizah, the Dead Sea Scrolls, and the Iron Age ostraca have gone through a digitization process, making most of the documents accessible in a digital format [6], [7], [8], though in varying degrees of accessibility and quality.

In cases of incised inscriptions (e.g., some of the Iron Age documents are incised on rocks or pottery sherds), RTI [9] may be

beneficial. The basic idea of RTI is that the camera and artifact are fixed, while the light angle varies over a spherical dome, creating a series of images with shading effects.

Another imaging technique used to improve the legibility of inscriptions is Infrared Reflectography (IRR). In IRR, broadband infrared light is emitted onto the surface of the artifact and the reflected photons are recorded. In cases where some traces of ink remain on the surface, perhaps invisible to the naked eye, IRR may help in improving the legibility. During the 1950s, IRR was used to document the Dead Sea Scrolls close to their time of discovery [8]. Since IRR improved the legibility of many parchment-based inscriptions, it was thought to be optimal for ostraca imaging as well. However, later studies showed that this is not always the case [10]. IRR has been used on the Cairo Genizah fragments, but mainly as part of material classification and not as a systematic way of improving legibility [10].

More recently, multi- and hyper-spectral imaging systems were introduced to the investigation and preservation of ancient documents. Similar to IRR, in these techniques, broadband light is emitted unto the artifact surface. However, here the reflected light is measured separately for different wavelengths; e.g., by utilizing narrow bandpass filters. That is, the imaging system counts the number of photons reflected in numerous wavelength ranges for each band. Multispectral imaging systems normally have a few dozens of recorded bands, whereas hyper-spectral systems have a few hundreds.

The study in [10] established a methodology for multispectral imaging optimized for ostraca. This technique has not only led to several new and improved readings of inscriptions [1], [10], [12], [13], [14], [15], but has also enabled the discovery of a brand-new inscription written on the back side of an already known ostracon, Arad 16 [1]; see Figure 3. The Dead Sea Scrolls have also been recorded systematically using multispectral imaging, and the images are available online [5].

Two less commonly used techniques are Raman Spectroscopy and micro-CT scanners. A Raman spectroscopy-based molecular scanner has been developed and used on ostraca in a proof-of-concept study, showcasing that ink traces can be mapped using

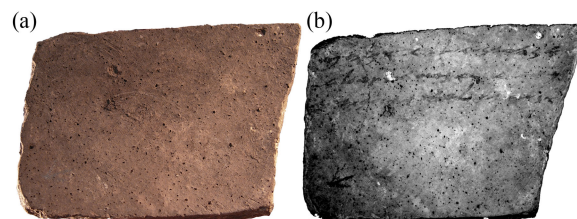


Figure 3

The backside of Arad 16 ostracon. (a) Standard digital image. (b) Multispectral band centered around 890 nm. Adapted from [1].

such instrumentation [16]. Micro-CT scans have been used to perform a successful digital unwarping and reading of an Engedi scroll (dated to the third or fourth century CE), avoiding the need for destructive procedures [17].

For a survey of similar and other applicable techniques in the setting of European documents (and objects of art), see [18]. Indeed, just as can be seen above for Hebrew documents, no universally applicable solution is offered, and all the methods are materials-dependent.

Image Analysis and Processing

Several image processing procedures can either be beneficial by themselves or as an intermediate stage for downstream processing tasks. The most common of such algorithms are binarization (i.e., creation of a black-and-white image of the inscription), and the closely-related segmentation (i.e., identifying and potentially “cropping” key areas in the image, such as characters). Both procedures can utilize grayscale, RGB, multi- or even hyper-spectral images as input data, with the former being the most common option.

Choosing the best input image is an interesting problem in its own right. Indeed, sometimes several images of the texts possessing different qualities are in existence, and an intelligent choice between them has to be made. To this end, several contrast measures have been suggested in the past, e.g., Weber, Michelson, RMS, and CMI (see their analysis and previous literature in [19]). However, the measured contrast of a given image can be misleading, as it might be affected by numerous potentially beneficial grayscale transformations performed by existing image processing software solutions. Therefore, the challenge is to measure the contrast of an image taking into account *all* its possible grayscale transformations. This problem leads to an alternative “Potential Contrast” (PC) measure, providing an analytic and extremely efficient solution. The PC measure was suggested in [19] and tested, among other types of images, on Iron Age ostraca (Horvat Radum, Horvat Uza) and one of the Dead Sea Scrolls.

Instead of choosing one particular image, an alternative approach would be to blend different images of the inscription into a single one, with multiple channels. Such a technique, combining new multispectral and old IRR images is offered in [20]. The proposed method consists of a two-step registration procedure—a coarse global transformation, followed by a fine local warping based on interest point matching. If the resulting “stacked” images are to be used directly by subsequent processing stages, these have to be adapted to accept multichannel images as their inputs.

Binarization is a common stage in computational paleography, mimicking the manual black and white facsimile creation performed by professional scholars. Several easy-to-implement

binarization algorithms are in widespread use. The most common global thresholding technique (setting one threshold for the whole grayscale image) is Otsu, maximizing the between-class variance. The common local techniques, making use of pixel values and their first and second-order statistics within a sliding window, are Bernsen, Niblack, and Sauvola (see previous literature and analysis in [21]).

Provided the special challenges of ancient degraded and noisy inscriptions, several specifically tailored binarization algorithms have been proposed. The multistep method in [22] begins with global thresholding, forming initial connected components (CCs). Within these, a quality evaluation is performed, resulting in pixels being assigned to either the foreground, the background, or transitional sets. First and second-order statistics of distance transforms of transitional pixels from the foreground are utilized to mark some characters as noisy. In such cases, a local morphological growing scheme that expands the characters to their final form is employed.

If a manual imperfect facsimile of the inscription exists, the algorithm in [21] suggests mining its information for binarization purposes. The process begins with preliminary global registration of the facsimile to the inscription image and continues with unconstrained elastic registration of each CC of the facsimile. A smoothing of the movements, akin to the median filter, is then performed to avoid local maxima. Subsequently, a binarization is conducted within a bounding octagon of each CC by setting the threshold according to the proportion of foreground pixels within the registered facsimile. Finally, an optional speckle-removing procedure can be performed. An example of binarization outputs for this and other algorithms, applied to Iron Age Arad 1 ostrakon, can be seen in Figure 4.

The binarizations can also be improved by various means. The method in [23] is based on a sparse dictionary-learning technique. A black and white dictionary is constructed from a clear source (such as a facsimile) and learned by k-medians, k-medoids, or extensive dictionary techniques. For each patch in the existing imperfect binarization, the most suitable replacement from the dictionary is chosen via a minimization procedure. The results of this study indicate that the k-medians and k-medoids methods are sound, with k-medians algorithm demonstrating better robustness to initial database shrinkage.

Alternative binarizations’ improvement solution, chosen in [24], is based on PixelCNN++, an autoregressive generative model, designed for image data. Four configurations are suggested: unmodified baseline; single model adaptive orientation; single model adaptive orientation, conditional; and multimodel adaptive orientation. Performance-wise, it seems that adaptive single- and multimodels achieve the best PSNR results.

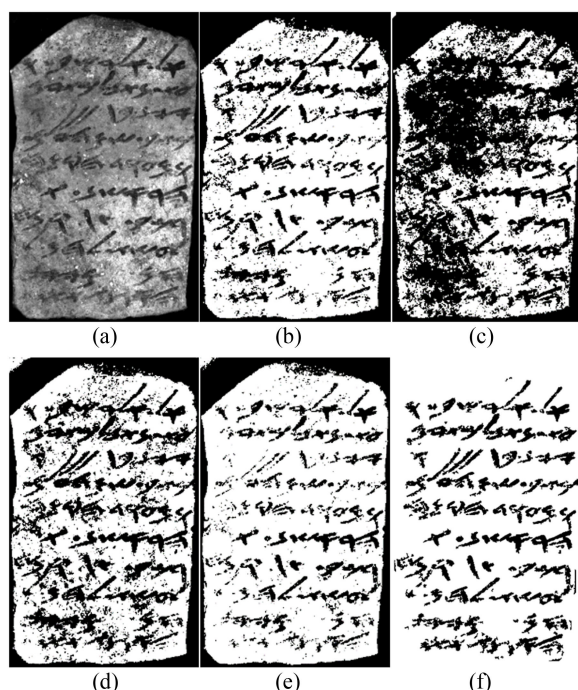


Figure 4

Comparison of binarizations for Arad 1 ostrakon.

(a) Inscription grayscale image. (b) Otsu. (c) Bernsen. (d) Niblack. (e) Sauvola. (f) Shaus et al. 2012. Adapted from [21].

Finding the characters' segmentation is a closely related problem. Indeed, upon achieving a satisfactory binarization, a segmentation might be as easy as extracting the CCs—perhaps several of them, if a complex character is of interest. A morphological filter taking such considerations into account is described in [22]. It consists of structuring element generation, character extraction, character validation, and structuring element adaptation stages.

However, a binarization might not be easily attainable, and one might be tempted to segment characters directly from the image itself. In such a case, a semi-automatic approach might be suggested, as in [25]. This study proposes a reconstruction of characters stroke-by-stroke, with a minimal user input. A stroke is defined as a piecewise-smooth part of a character with a specific radius of writing at each point, resulting from the act of writing. This definition reflects and reconstructs scribe's read pen movements. An energy functional minimization procedure is employed in order to find a solution to the corresponding optimization problem, resulting in plausible results—some of which can be seen in Figure 5.

To sum up, just as in the case of image acquisition, image analysis and processing are to a large extent domain-dependent. Indeed, this is also the case for various types of European texts, as reflected in the survey [26]. In fact, despite existing

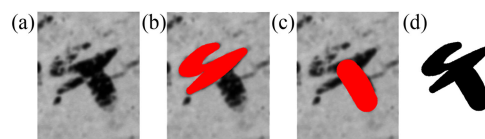


Figure 5

Restoration of a character waw from Arad ostrakon 24.

(a) Original image. (b) and (c) reconstructed strokes. (d) Resulting character restoration. Adapted from [25].

benchmark databases for such documents, there continues to be a debate on the proper way to evaluate the algorithms, given that humans do not always agree on the ground truth at the pixel level (see discussion in [27]). Moreover, [26] also mentions the foremost issue of algorithm reproducibility: it is not always clear which parts (pre-/postprocessing, parameter tuning, local threshold selection) of previous algorithms were successful and should be reused; few authors provide a full documentation and make their code publicly available. Generalizing algorithms across domains also remains a challenge. Finally, it was noted that methods tend to break when there are large stains, or in the presence of border noise.

Handwriting Analysis

Handwriting is considered to be a unique “fingerprint” that characterizes a scribe. The distinct style of writing plays a significant role in identifying writers, tracking the evolution of the script, and dating the inscriptions. Although classical paleography aims at answering these questions, computational handwriting analysis can supplement the traditional studies by providing efficient and statistically justified evidence, shedding light on long-debated historical questions.

Among the main challenges of ancient handwriting analysis are the limited number of available documents (i.e., this is a case of “small” rather than “big data”), the lack of labeled reference data, as well as the poor preservation level of the documents. Due to these complexities, which vary across corpora, new analytical methods, not necessarily based on deep learning, had to be developed.

The common handwriting analysis tasks are as follows.

- 1) writers' comparison or classification;
- 2) fragments' matching (finding 'joins');
- 3) letters' prior estimation;
- 4) dating based on handwriting style.

Although the medium and the research question may vary, there is a common flow that can be observed in most of the studies that deal with handwriting analysis. In Figure 6, we illustrate this flow, used in [28], [29], [30], and [31].

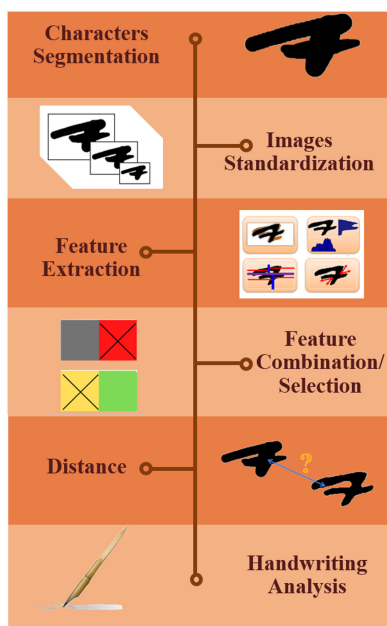


Figure 6

Common stages of computational handwriting analysis.

Typically, document analysis focuses on analyzing the geometrical shapes of the characters in the text, however, at times, other “global” geometrical aspects of the document are considered (see [32]). Table 1 summarizes the common steps performed by key papers in the field, representing, to the best of our knowledge and judgment, the culmination of efforts of all the existing research groups in the field.

Writers’ Comparison or Classification

Commonly, the methods are employed on a character level. Therefore, initial data cleansing and standardization procedures, e.g., rotating the characters according to their line inclination [28], [29], aligning them to one another, and resizing them based on a predefined scale, are of utmost importance.

Once the character images are standardized, a set of features, incorporating their various geometrical aspects, is extracted. There is a wide variety of possible features that may be useful for handwriting analysis. The decision regarding the optimal set of features might be influenced by the research question at hand, the information in use (i.e., the original images, binary images, or multispectral cube of the inscriptions), the amount of available data, and its variability. Usually, due to the degradation of historical documents, automatic handwriting analysis tasks rely upon features extracted from clear images of the characters (i.e., characters’ binarizations, see the previous section for additional details).

The extracted features can be adapted from the domains of Computer Vision and Document Analysis. For example, in [22],

the authors designed shape-related features that are sensitive to small differences in the character’s forms, while taking into account the various aspects of the proportion between the letter volume and its background. In the studies reported in [30], [33], and [34], several features were utilized, referring to aspects such as the character’s overall shape, the angles between strokes, the character’s center of gravity, as well as its horizontal and vertical projections (the full list of features consists of SIFT, Zernike, DCT, Kd-tree, image projections, L1 and CMI (see [30] for additional information)).

In [29], the authors also used SIFT for the comparison task, while in [31], three-level feature extraction was employed: (a) texture-level captured the curvature and slant of the contours of characters using the Hinge method; (b) character-shape (allograph) level, namely Fraglets, using a Neural Network that reduces their dimension; (c) adjoined feature, combining the previous two.

Later, the extracted features can be combined to provide an embedding into the features’ domain. This task can be addressed either by *feature combination* [29], [30] or by *feature selection* [22], [31]. Typically, the feature vectors are normalized prior to performing either activity, in order to deal with features on the same scale [29], [30], [31].

In the case of feature combination, all the features are blended into a single descriptor (e.g., a vector), retaining all the meaningful information. In [30], [33], and [34], the features were combined into a single, generalized feature vector, that described each character by the degree of its proximity to all of the characters. In [29], the features were combined using the bag-of-visual-keyword method, constructing a histogram based on a predefined dictionary. In [31], PCA was applied, while in [22], Fisher Linear Discriminant Analysis was used.

The feature selection approach assumes that some features might not have a disproportionately large influence on the result, and therefore, selecting them (and disregarding the others) can be beneficial. For example, in [19], the sequential forward floating selection was utilized.

Using the embedded representation of the characters, the degree of similarity between documents can be measured according to combined features, representing various aspects of the text. Thus, distances between elements are measured and statistical inference is performed in order to answer the research questions posed. The distances can be simple Euclidian distance (e.g., [29], [30], [33], [34]), or a distance tailored to the descriptors of the feature (e.g., Pearson’s chi-square test statistics in [31]).

Finally, the document analysis task can be addressed. As mentioned above, several goals can be tackled. In [22], the authors dealt with the writer identification in Medieval Hebrew calligraphic documents, by utilizing k-nearest neighbors, with $k =$

TABLE 1. Summary of Writer Identification Performed in Several Key Papers

Paper	Corpus	Task	Feature Extraction	Feature Combination / Selection	Handwriting Analysis Method
Bar Yosef et al. [22]	Medieval calligraphic documents dated to 14th–16th CE.	Writer identification	Different aspects of letter area with respect to its background	Sequential forward floating selection, Fisher linear discriminant analysis	k-nearest neighbors & Linear Bayes classifier
Wolf et al. [29]	Cairo Genizah, written mainly in the 10th–15th CE	Handwriting matching and classification	SIFT	Bag-of-visual-keyword	Nearest neighbor
Faigenbaum-Golovin et al. [30], [33], Shaus et al. [34]	Arad and Samaria ostraca dated ca. 600 BCE and 8th century BCE.	Estimation of the number of scribes in a given corpus	Template matching, SIFT, KD-tree, horizontal and vertical projections, Zernike moments, DCT.	Each character is described by its distance from all other characters of the same type	Hypothesis testing
Shaus and Turkel [35], Shaus et al. [34]	Arad ostraca dated ca. 600 BCE and 8th century BCE.	Estimation of the number of scribes in a given corpus	Binary pixel patterns— 3×3 black and white patches	Each character is described by a histogram of patches	Hypothesis testing
Popović et al. [31]	Dead Sea Scrolls, dated to 4th century BCE - 2nd century CE, Great Isaiah Scroll.	Writer identification	Curvature and slant of the contours, characters shape level using Neural Networks	PCA	Manual visual analysis

1, ..., 5, or a Linear Bayes classifier. Their results indicate good results when the classification was based on the majority rule. In [29], the writer attribution method was applied to texts originating from the Cairo Genizah. The proposed method selected the nearest neighbor from reference data. The reported accuracy was acceptable for handwriting recognition, with moderately good classification results.

In [30], estimation of the minimal number of writers is addressed by the document comparison method, and by testing the null hypothesis “both texts were written by the same author.” The p -values obtained for each relevant letter in the alphabet are later combined into a single p -value via the well-established Fisher’s method. The *minimal* number of writers was established by finding the largest pairwise distinct document set (“clique”). When evaluated on Modern Hebrew documents, the algorithm yielded 2% of False Positive and 2% of False Negative results. Applied to the corpus

of 18 texts from the tiny Arad fortress located in the Judean Desert (ca. 600 BCE), four distinct scribes were identified, with repercussions regarding the widespread literacy in the late monarchic Judah. Further upgrades of the method in [34] yielded five distinct scribes in the same setting.

In another related study [33], the Arad corpus was utilized to establish the confusion matrices for inscriptions of different lengths. These formed the foundation for an enhanced toolbox providing the *maximum likelihood estimate* of the number of hands in another corpus under investigation. The method was employed on 31 ostraca from the 8th century BCE Samaria, the capital of the biblical kingdom of Israel. It was established that these texts were most likely written by only two contemporaneous scribes recording the shipments in Samaria over the span of seven years. These outcomes contrast with the results of relatively widespread literacy in Arad mentioned above.

The study in [34] computationally analyzed the Arad ostraca handwriting; the results were contrasted with the examination by forensic document expert. The study demonstrated substantial agreement between the results of these independent methods of investigation. Remarkably, the forensic examination revealed a high probability of at least 12 writers within the analyzed corpus, indicating that the computational methods are rather conservative in their conclusions.

We also note two additional lines of research, published in [34], [35], and [36], which performed writer identification on the Arad corpus, testing the null hypothesis that two given inscriptions were written by the same author. The study in [35] utilized a histogram of binary pixel patterns (3×3 black and white patches) to represent the characters, and performed multiple experiments of the Kolmogorov-Smirnov test for each letter and each patch, combining the resulting p -values via Fisher's method. Some corrections taking potential correlations between p -values were introduced in [34]. The different method suggested in [36] projected letters' binarizations via multidimensional scaling and performed a two-sample t -test.

In [31], the authors examined the Great Isaiah Scroll, belonging to the corpus of Dead Sea Scrolls. The authors performed a manual analysis of the visualized data embedded in 3-D via PCA. Unfortunately, no accuracy estimation on ground truth data was provided within the scope of that paper.

Fragments' Matching

Another interesting and related procedure is the matching of fragments. Finding "joins" between different pieces of the same document, as presented in [32], avoids the weighty pipeline required and used by handwriting analysis, and substitutes it for a global analysis. Several measurements quantifying the geometrical aspects of the document are considered: number of lines, average line height, the standard deviation of line height, average space between lines, the standard deviation of interline space, and different bounding box aspects.

Letters' Prior Estimation

Yet another handwriting analysis task is the "letter prior" estimation. In paleography, it is oftentimes assumed that each writer has a prototype for each letter in the alphabet. In large paleographical studies, a manually created paleographic table containing such representative characters is often provided. Such tables can be used for style and chronological analysis. The laborious process of priors' estimation is a natural candidate for computational implementation. Indeed, several prior estimation methods were proposed in the past. In [28], the priors were extracted using fragments of characters that were registered via different techniques. In [37], median calculation per pixel was performed, while in [36], a histogram of representative medoids was constructed.

Dating Based on Handwriting Style

Another task, important for historical reconstructions, is the period attribution ("dating") based on computerized handwriting style analysis. The only example of handling this question in our context, provided in [20], tested the discrimination capabilities of several measures: interest-point displacement norm, Jaccard index, and mutual information. On Dead Sea Scrolls, which provided material from Herodian, Hasmonean, and Hellenistic-Roman inscriptions, the displacement norm and mutual information managed to discriminate between the Hellenistic-Roman period and each of the other two, whereas the Jaccard index was only able to discriminate between the Hellenistic-Roman and Hasmonean periods. For comparison, see a survey dealing with dating European texts [38]; methods incorporating linguistic features (a challenge for ancient Hebrew corpora) are found to be especially beneficial.

Summary

In this section, we reviewed several methods for performing handwriting analysis tasks. A decade or so ago, when the field of computerized paleography was in its infancy, no off-the shelf methods existed, and therefore brand-new algorithms had to be developed. Although, superficially, some of the research questions resembled each other, in fact, the writing medium, the script, the documents' level of preservation, as well as the amounts of available data, posed unique challenges and dictated distinct solutions. Implementing these often (and regrettably) relied on shortcuts, such as assumptions regarding the input data, or insufficiently robust and evaluated pipelines. E.g., several methods assumed an existence of noise-free binarization (this is a pre-processing step in [22], [30], and [33]); some (such as [30], [33], and [34]) utilized segmented and labeled characters. Certain algorithms (e.g., [27]) relied on a single feature for the comparison task,—this may have not necessarily captured all the aspects of the handwriting and could have resulted in a nonrobust result. Accuracy evaluation is a highly advisable step, especially when a new method is suggested, however, such stage was sometimes omitted (as in [31]).

Further Tasks

Additional classical paleographic questions can be answered by computational means. Examples include a comparison of manually created facsimiles, transcription-assistance tools, OCR, and transcripts alignment.

As already mentioned, the discipline of paleography relies heavily on manually-drawn facsimiles of ancient inscriptions. This practice may unintentionally mix up documentation and interpretation. Hence, evaluating the quality of the facsimile is important. In a way, such an evaluation might be seen as

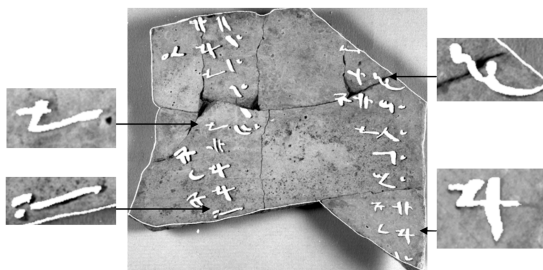


Figure 7

Illustration of facsimile quality-assessment procedure for Arad ostracon 34, adapted from [39].

complementary to the comparison of the input images (see above). Indeed, the toolkit in-use in [39] is also reliant on the registration of the facsimile to the inscription's image, and the CMI metric. The results of the study showcase the amount of bias present in the facsimiles produced by different scholars. For an example, see Figure 7; note the shadows indicating an imperfect match between the inscription's image and the manual facsimile.

On the other hand, Faigenbaum et al. [40] suggest that facsimiles can also be evaluated as-is on an individual glyph (character or ligature) level. This can be done via intrinsic quality assessment features such as stroke width consistency, presence of small CCs (stains), edge noise, and the average edge curvature. Linear and tree-based combinations of these features are also considered. The new methodology is tested and shown to be nearly as sound as human judgment.

The creation of transcription is another crucial task. Provided the inscriptions' state preservation, with missing, effaced, and dubious characters, searching for an appropriate reading might be extremely challenging and tedious. This is the case even if some very detailed dictionary of probably occurring words is employed. For ancient Hebrew, such problems can be handled by the tool presented in [41]. It combines the user's input (specific certain characters, possible characters, range of numbers of characters in specific words, etc.) with computationally efficient search in pre-existing dictionaries via a regular expressions engine, available online. For an example see Figure 8.

In the case of pre-existing transcriptions, it is sometimes desirable to align them to the inscription's image. This task is handled in [42], where the suggested pipeline includes baseline segmentation, line polygon extraction, automated transcription via a hybrid CNN-RNN method, and alignment of Unicode characters in transcription with the characters in the image through either optical SIFT-flow or with OCR results.

The OCR of ancient Hebrew text is dealt with, *inter alia* and rather briefly in [22] and [35]. In the former case, a solution

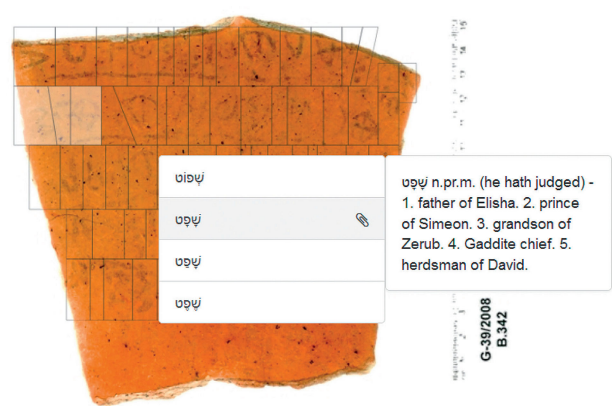


Figure 8

Regular-expression driven dictionary search in Qeiyafa ostracon, adapted from [41].

based on morphological analysis is employed, while the latter study uses a CNN-RNN hybrid, trained in a supervised manner via a turn-key Kraken engine, with default parameters. The performance on large-scale test cases, especially in [35], is not particularly remarkable—as expected for highly degraded ancient texts with low amounts of training material. The survey [43] concurs that the OCR for ancient (non-Hebrew) texts is a challenging endeavor, and in general not much work has been done for such documents.

Conclusion

We conclude this article by discussing the applicability of the surveyed methods for other ancient scripts and highlighting possible future research directions. Herein we omit the discussion pertaining to modern scripts, since the technologies available for both online and offline analysis of these types of clean data are rather advanced and beyond the scope of the current article.

The diverse algorithms presented in this article handled ancient Hebrew inscriptions from different historical periods. Despite certain evolution in the alphabet, some characteristics of it had persisted throughout the ages: the number of its letters is 22, and in fact, the letters themselves (*alef*, *bet*, etc.) remain the same—though their shape certainly changed; most of the characters are disjoint and “touching letters” are quite rare (unlike cursive writings in other types of writing, such as modern Latin scripts or Arabic); most of the characters are composed of just a handful of strokes. Naturally, as discussed above, all these ancient corpora contain different types of noise and degradation.

Therefore, it is reasonable to assume, that for a script possessing similar characteristics, such as the closely related

ancient Phoenician/Punic writing, or the earliest variants of Greek and Latin script, many of the approaches described in this article would be almost immediately applicable.

For rather different languages and scripts, the applicability would require more complicated adaptations. For some ancient writing systems, the number of symbols in use might be significantly larger, e.g., in ancient classical Egyptian writing with hundreds of common signs, or East Asian scripts with thousands of symbols—and possibly a large number of strokes. It is quite likely, that for such scripts, a substantial modification effort would be required, the amount of input material would have to be larger, and in some instances, different image acquisition and image processing techniques would have to be developed (e.g., for cuneiform writing, it might be a high-resolution development of RTI combined with a tailored image blending procedure). In the case of purely cursive languages, segmenting the characters is a known challenge not covered herein; perhaps if common words or ligatures are present, these can be employed in handwriting comparisons, OCR, etc.

Some possible future research directions might be considered. In image acquisition, these include wider employment of hyperspectral imaging, a utilization of UV and Far-IR spectral information, macro XRF (as in [44]), and Raman-based imaging (suggested in [16]).

In the context of information processing, multimodal data frequently refers to audio and video signals recording the same scene, thus containing complementary information. In case of various imaging techniques recording the same object (i.e., written artifacts), the complementary information obtained relates to different physical properties. It is expected that fusing such data would encompass completely different approaches than the ones used in the audio-video problem of multimodal signal processing problems. Another promising avenue of research is a fusion of different kinds of information, e.g., an underlying language model may benefit text segmentation.

Another open problem that could ease the process of analyzing handwriting significantly is devising a fully automatic binarization framework. Though this problem seems to be solved in standard document analysis, the issue with ancient inscriptions is that in many cases some of the degradations have the properties of natural signals or missing data. Thus, classical approaches fail time and time again. It is expected that modern machine learning approaches coupled with careful modeling of the signals should break the ground in this respect.

In the recent decade, deep neural networks (DNNs) have advanced the performance of machine learning significantly in a diverse set of tasks. Therefore, the fact that we do not report in this survey any DNN-based approach is somewhat surprising. However, normally, training DNNs requires a substantial number of samples—an unattainable scenario in the context of ancient Hebrew inscriptions. Furthermore, using

artificially generated data, or data generated from other sources or tasks (e.g., as in transfer learning), may introduce biases that are hard to account for. Nonetheless, we expect that in the near future modern self-supervised approaches could be designed to deal with such issues. In fact, such an approach has already been employed in the case of Vatican historical materials [45].

Additionally, in handwriting analysis, estimating the most likely number of scribes in a corpus with no reliance on external data might be an interesting challenge. An even more challenging and intriguing direction is the characterization of joint information that is common to a school of scribes. Being able to define and quantify this in a proper manner would open the gates to new research questions that can be tackled.

Finally, in our view, to this day, two crucial topics did not receive enough scholarly attention. The first is the need to engage the paleographic community by providing and supporting easily usable and deployable software toolkits. The second is the need to invest more computational and experimental effort to devise best conservation practices for ancient inscriptions, allowing future generations to experience and enjoy the remarkable materials currently entrusted to us.

Acknowledgment

This work was supported in part by the ISF under Grant 2062/18 and in part by Jacques Chahine (made through the French Friends of Tel Aviv University). The work of Shira Faigenbaum-Golovin was supported by the Eric and Wendy Schmidt Fund for Strategic Innovation, and the Zuckerman-CHE STEM Program, by Duke University, and by the Simons Foundation under Grant Math+X 400837. The authors would like to thank Israel Finkelstein, David Levin, Eli Piasetzky, Eli Turkel, Michael Cordonsky, Yana Gerber, Eythan Levy, and Anat Mendel-Geberovich for their kind guidance, cooperation, and assistance over the years. The authors would also like to thank the *BITS* editor and the anonymous reviewers for their helpful comments and suggestions. Shira Faigenbaum-Golovin, Arie Shaus, and Barak Sober contributed equally to this work.

References

- [1] S. Faigenbaum-Golovin et al., “Multispectral imaging reveals biblical-period inscription unnoticed for half a century,” *PLoS One*, vol. 12, no. 6, 2017, Art. no. e0178400, doi: [10.1371/journal.pone.0178400](https://doi.org/10.1371/journal.pone.0178400).
- [2] S. Faigenbaum-Golovin, C. A. Rollston, E. Piasetzky, B. Sober, and I. Finkelstein, “The Ophel (Jerusalem) ostrakon in light of new multispectral images,” *Semitica*, vol. 57, pp. 113–137, 2015, doi: [10.2143/SE.57.0.3115458](https://doi.org/10.2143/SE.57.0.3115458).

- [3] Fragment of a Torah Scroll, 13th century?, found in the Cairo Genizah (MS. Heb. a. 4). Accessed: Aug. 2022. [Online]. Available: [https://commons.wikimedia.org/wiki/File:Torah_Scroll_\(MS_heb_a_4\).jpg](https://commons.wikimedia.org/wiki/File:Torah_Scroll_(MS_heb_a_4).jpg)
- [4] Dead Sea Scroll fragment 4Q7, Genesis 1. Accessed: Aug. 2022. [Online]. Available: [https://commons.wikimedia.org/wiki/File:Genesis_1_Dead_Sea_Scroll_\(Cropped\).jpg](https://commons.wikimedia.org/wiki/File:Genesis_1_Dead_Sea_Scroll_(Cropped).jpg)
- [5] The Leon Levy Dead Sea Scrolls Digital Library. Accessed: Aug. 2022. [Online]. Available: <https://www.deadseascrolls.org.il>
- [6] L. Hunt, M. Lundberg, and B. Zuckerman, "InscriptiFact: A virtual archive of ancient inscriptions from the near east," *Int. J. Dig. Libraries*, vol. 5, no. 3, pp. 153–166, 2005, doi: [10.1007/s00799-004-0102-z](https://doi.org/10.1007/s00799-004-0102-z).
- [7] Y. Choueka, "Computerizing the Cairo Genizah: Aims, methodologies and achievements," *Ginzei Qedem*, vol. 8, pp. 9–30, 2012.
- [8] P. Shor, M. Manfredi, G. H. Bearman, E. Marengo, K. Boydston, and W. A. Christens-Barry, "The Leon Levy Dead Sea Scrolls digital library: The digitization project of the Dead Sea Scrolls," *J. Eastern Mediterranean Archaeol. Heritage Stud.*, vol. 2, no. 2, pp. 71–89, 2014, doi: [10.5325/jeasmedarcherstu.2.2.0071](https://doi.org/10.5325/jeasmedarcherstu.2.2.0071).
- [9] R. Woodham, "Photometric method for determining surface orientation from multiple images," *Opt. Eng.*, vol. 19, no. 1, pp. 139–144, 1980, doi: [10.1117/12.7972479](https://doi.org/10.1117/12.7972479).
- [10] S. Faigenbaum et al., "Multispectral images of ostraca: Acquisition and analysis," *J. Archaeological Sci.*, vol. 39, no. 12, pp. 3581–3590, 2012, doi: [10.1016/j.jas.2012.06.013](https://doi.org/10.1016/j.jas.2012.06.013).
- [11] Z. Cohen, "Experimental methods," in *Composition Analysis of Writing Materials in Cairo Genizah Documents* (Cambridge Genizah Studies Series 15). Leiden, The Netherlands: Brill, 2021, pp. 53–66, doi: [10.1163/97890004469358_005](https://doi.org/10.1163/97890004469358_005).
- [12] B. Sober et al., "Multispectral imaging as a tool for enhancing the reading of ostraca," *Palestine Exploration Quart.*, vol. 146, no. 3, pp. 185–197, 2014, doi: [10.1179/0031032814Z.000000000101](https://doi.org/10.1179/0031032814Z.000000000101).
- [13] S. Faigenbaum, B. Sober, M. Moinester, and E. Piasetzky, "Multispectral imaging of Tel Malhata ostraca," in *Tel Malhata, A Central City in the Biblical Negev* (Institute of Archaeology, Tel Aviv University Monograph Series 32), I. Beit-Arie and L. Freud, Eds., Tel Aviv, Israel: Eisenbrauns, 2015, pp. 510–513, doi: [10.5325/j.ctv1bxh406.11](https://doi.org/10.5325/j.ctv1bxh406.11).
- [14] A. Mendel-Geberovich et al., "A renewed reading of Hebrew ostraca from cave A-2 at Ramat Beit Shemesh (Nahal Yarmut), based on multispectral imaging," *Vetus Testamentum*, vol. 69, no. 4/5, pp. 682–701, 2019, doi: [10.1163/15685330-00001370](https://doi.org/10.1163/15685330-00001370).
- [15] S. Faigenbaum-Golovin, I. Finkelstein, E. Levy, N. Na'aman, and E. Piasetzky, "Arad ostraca 24 side A," *Semitica*, vol. 62, pp. 43–68, 2020, doi: [10.2143/SE.62.0.3288851](https://doi.org/10.2143/SE.62.0.3288851).
- [16] A. Shaus et al., "Raman binary mapping of Iron Age ostraca in unknown material composition and high fluorescence setting – A proof of concept," *Archaeometry*, vol. 61, no. 2, pp. 459–469, 2019, doi: [10.1111/arcm.12419](https://doi.org/10.1111/arcm.12419).
- [17] W. B. Seales, C. S. Parker, M. Segal, E. Tov, P. Shor, and Y. Porath, "From damage to discovery via virtual unwrapping: Reading the scroll from En-Gedi," *Sci. Adv.*, vol. 2, no. 9, 2016, Art. no. e1601247, doi: [10.1126/sciadv.1601247](https://doi.org/10.1126/sciadv.1601247).
- [18] A. Tonazzini et al., "Analytical and mathematical methods for revealing hidden details in ancient manuscripts and paintings: A review," *J. Adv. Res.*, vol. 17, pp. 31–42, 2019, doi: [10.1016/j.jare.2019.01.003](https://doi.org/10.1016/j.jare.2019.01.003).
- [19] S. F.-G. A. Shaus, B. Sober, E. Turkel, and E. Piasetzky, "Potential contrast - A new image quality measure," in *Proc. IS&T Int. Symp. Electron. Imag., Image Qual. System Perform. XIV Conf.*, 2017, pp. 52–58, doi: [10.2352/ISSN.2470-1173.2017.12.IQSP-226](https://doi.org/10.2352/ISSN.2470-1173.2017.12.IQSP-226).
- [20] Y. Zarai, T. Lavee, N. Dershowitz, and L. Wolf, "Integrating copies obtained from old and new preservation efforts," in *Proc. 12th Int. Conf. Document Anal. Recognit.*, 2013, pp. 47–51, doi: [10.1109/ICDAR.2013.18](https://doi.org/10.1109/ICDAR.2013.18).
- [21] A. Shaus, E. Turkel, and E. Piasetzky, "Binarization of first temple period inscriptions: Performance of existing algorithms and a new registration based scheme," in *Proc. 13th Int. Conf. Front. Handwriting Recognit.*, 2012, pp. 645–650, doi: [10.1109/ICFHR.2012.187](https://doi.org/10.1109/ICFHR.2012.187).
- [22] I. Bar Yosef, I. Beckman, K. Kedem, and I. Dinstein, "Binarization, character extraction, and writer identification of historical Hebrew calligraphy documents," *Int. J. Document Anal. Recognit.*, vol. 9, no. 2–4, pp. 89–99, 2007, doi: [10.1007/s10032-007-0041-5](https://doi.org/10.1007/s10032-007-0041-5).
- [23] A. Shaus, B. Sober, E. Turkel, and E. Piasetzky, "Improving binarization via sparse methods," in *Proc. 16th Int. Graphonomics Soc. Conf.*, 2013, pp. 163–166.
- [24] L. Uzan, N. Dershowitz, and L. Wolf, "Qumran letter restoration by rotation and reflection modified PixelCNN," in *Proc. 14th IAPR Int. Conf. Document Anal. Recognit.*, 2017, pp. 23–29, doi: [10.1109/ICDAR.2017.14](https://doi.org/10.1109/ICDAR.2017.14).
- [25] B. Sober and D. Levin, "Computer aided restoration of handwritten character strokes," *Comput. Aided Des.*, vol. 89, pp. 12–24, 2017, doi: [10.1016/j.cad.2017.04.005](https://doi.org/10.1016/j.cad.2017.04.005).
- [26] C. Tensmeyer and T. Martinez, "Historical document image binarization: A review," *SN Comput. Sci.*, vol. 1, 2020, Art. no. 173, doi: [10.1007/s42979-020-00176-1](https://doi.org/10.1007/s42979-020-00176-1).
- [27] A. Shaus, B. Sober, E. Turkel, and E. Piasetzky, "Beyond the ground truth: Alternative quality measures of document binarizations," in *Proc. 15th Int. Conf. Front. Handwriting Recognit.*, 2016, pp. 495–500, doi: [10.1109/ICFHR.2016.0097](https://doi.org/10.1109/ICFHR.2016.0097).
- [28] I. Bar Yosef, A. Mokeichev, K. Kedem, I. Dinstein, and U. Ehrlich, "Adaptive shape prior for recognition and variational segmentation of degraded historical characters," *Pattern Recognit.*, vol. 42, no. 12, pp. 3348–3354, 2009, doi: [10.1016/j.patcog.2008.10.005](https://doi.org/10.1016/j.patcog.2008.10.005).
- [29] L. Wolf, N. Dershowitz, L. Potikha, T. German, R. Shweka, and Y. Choueka, "Automatic paleographic exploration of Genizah manuscripts," in *Codicology and Palaeography in the Digital Age 2*, F. Fischer, C. Fritze, and G. Vogeler, Eds., Norderstedt, Germany: Books on Demand, pp. 157–179, 2011.

- [30] S. Faigenbaum-Golovin et al., "Algorithmic handwriting analysis of Judah's military correspondence sheds light on composition of biblical texts," *Proc. Nat. Acad. Sci. USA*, vol. 113, no. 17, pp. 4664–4669, 2016, doi: [10.1073/pnas.1522200113](https://doi.org/10.1073/pnas.1522200113).
- [31] M. Popović, M. A. Dhali, and L. Schomaker, "Artificial intelligence based writer identification generates new evidence for the unknown scribes of the dead sea scrolls exemplified by the Great Isaiah Scroll (1QIsaa)," *PLoS One*, vol. 16, no. 4, 2021, Art. no. e0249769, doi: [10.1371/journal.pone.0249769](https://doi.org/10.1371/journal.pone.0249769).
- [32] L. Wolf et al., "Identifying join candidates in the Cairo Genizah," *Int. J. Comput. Vis.*, vol. 94, no. 1, pp. 118–135, 2011, doi: [10.1007/s11263-010-0389-8](https://doi.org/10.1007/s11263-010-0389-8).
- [33] S. Faigenbaum-Golovin, A. Shaus, B. Sober, E. Turkel, E. Piasetzky, and I. Finkelstein, "Algorithmic handwriting analysis of the Samaria inscriptions illuminates bureaucratic apparatus in biblical Israel," *PLoS One*, vol. 15, no. 1, 2020, Art. no. e0227452, doi: [10.1371/journal.pone.0227452](https://doi.org/10.1371/journal.pone.0227452).
- [34] A. Shaus, Y. Gerber, S. Faigenbaum-Golovin, B. Sober, E. Piasetzky, and I. Finkelstein, "Forensic document examination and algorithmic handwriting analysis of Judahite biblical period inscriptions reveal significant literacy level," *PLoS One*, vol. 15, no. 9, 2020, Art. no. e0237962, doi: [10.1371/journal.pone.0237962](https://doi.org/10.1371/journal.pone.0237962).
- [35] A. Shaus and E. Turkel, "Writer identification in modern and historical documents via binary pixel patterns, Kolmogorov-Smirnov test and Fisher's method," *J. Imag. Sci. Technol.*, vol. 61, no. 1, pp. 010404–1–010404–9, 2017, doi: [10.2352/J.ImagingSci.Technol.2017.61.1.010404](https://doi.org/10.2352/J.ImagingSci.Technol.2017.61.1.010404).
- [36] S. Faigenbaum-Golovin, D. Levin, E. Piasetzky, and I. Finkelstein, "Writer characterization and identification in short modern and historical documents: Reconsidering paleographic tables," in *Proc. 19th ACM Symp. Document Eng.*, 2019, Art. no. 23, doi: [10.1145/3342558.3345413](https://doi.org/10.1145/3342558.3345413).
- [37] A. Shaus and E. Turkel, "Towards letter shape prior and paleographic tables estimation in Hebrew First Temple period ostraca," in *Proc. 4th Int. Workshop Historical Document Imag. Process.*, 2017, pp. 13–18, doi: [10.1145/3151509.3151511](https://doi.org/10.1145/3151509.3151511).
- [38] S. Boldsen and F. Wahlberg, "Survey and reproduction of computational approaches to dating of historical texts," in *Proc. 23rd Nordic Conf. Comput. Linguistics*, 2021, pp. 145–156.
- [39] A. Shaus, E. Turkel, and E. Piasetzky, "Quality evaluation of facsimiles of Hebrew First Temple period inscriptions," in *Proc. 10th IAPR Int. Workshop Document Anal. Syst.*, 2012, pp. 170–174, doi: [10.1109/DAS.2012.70](https://doi.org/10.1109/DAS.2012.70).
- [40] S. Faigenbaum, A. Shaus, B. Sober, E. Turkel, and E. Piasetzky, "Evaluating glyph binarizations based on their properties," in *Proc. ACM Symp. Document Eng.*, 2013, pp. 127–130, doi: [10.1145/2494266.2494318](https://doi.org/10.1145/2494266.2494318).
- [41] E. Levy and F. Pluquet, "Computer experiments on the Khirbet Qeiyafa ostrakon," *Dig. Scholarship Humanities*, vol. 32, no. 4, pp. 816–836, 2017, doi: [10.1093/lc/fqw028](https://doi.org/10.1093/lc/fqw028).
- [42] D. Stökl Ben Ezra, B. Brown-DeVost, N. Dershowitz, A. Pechorin, and B. Kiessling, "Transcription alignment for highly fragmentary historical manuscripts: The Dead Sea Scrolls," in *Proc. 17th Int. Conf. Front. Handwriting Recognit.*, 2020, pp. 361–366, doi: [10.1109/ICFHR2020.2020.00072](https://doi.org/10.1109/ICFHR2020.2020.00072).
- [43] S. R. Narang, M. K. Jindal, and M. Kumar, "Ancient text recognition: A review," *Artif. Intell. Rev.*, vol. 53, pp. 5517–5558, 2020, doi: [10.1007/s10462-020-09827-4](https://doi.org/10.1007/s10462-020-09827-4).
- [44] J. Powers et al., "X-ray fluorescence recovers writing from ancient inscriptions," *Zeitschrift für Papyrologie und Epigraphik*, vol. 152, pp. 221–227, 2015.
- [45] L. Lastilla, S. Ammirati, D. Firmani, N. Komodakis, P. Merialdo, and S. Scardapane, "Self-supervised learning for medieval handwriting identification: A case study from the Vatican apostolic library," *Inf. Process. Manag.*, vol. 59, no. 3, 2022, Art. no. 102875, doi: [10.1016/j.ipm.2022.102875](https://doi.org/10.1016/j.ipm.2022.102875).

Shira Faigenbaum-Golovin received the B.Sc., M.Sc., and Ph.D. degrees in applied mathematics from Tel-Aviv University, Tel Aviv-Yafo, Israel, in 2006, 2014, and 2021, respectively. She is currently a Phillip Griffiths Assistant Research Professor with Mathematics Department, Duke University, Durham, NC, USA, as well as with the Rhodes Interdisciplinary Initiative. She works on theory and applications in computer vision. For the last decade, she has studied historical documents via statistical analysis, pattern recognition, and machine learning tools. Her research interests include approximation in high-dimensional space, manifold learning, as well as addressing various challenges posed by high-dimensional data.

Arie Shaus received the B.Sc. degree in mathematics (*summa cum laude*), the B.A. degree in archaeology and ancient Near Eastern cultures (*summa cum laude*), the M.Sc. degree in applied mathematics (*magna cum laude*), and the Ph.D. degree in applied mathematics, from Tel Aviv University, Tel Aviv, Israel, in 2006, 2011, 2015, and 2017, respectively. Currently splits his time between the Harvard Medical School, and the Department of Archaeology, Tel Aviv University, Tel Aviv-Yafo, Israel. His research interests include image processing, machine learning, biblical archaeology, and ancient DNA.

Barak Sober received the B.Sc. degree in mathematics and philosophy and M.Sc. and Ph.D. degrees in applied mathematics from Tel Aviv University in 2009, 2014, and 2019, respectively. He is currently an Assistant Professor of statistics and data science and digital humanities with the Hebrew University of Jerusalem, Jerusalem, Israel. His research interests include analysis of high dimensional data from a geometrical perspective and the application of mathematical modeling and statistical methods in the arts and humanities with the aim of performing empirically based scientific investigations in these fields.