

Digitizing and parsing semi-structured historical administrative documents from the G.I. Bill mortgage guarantee program

Digitizing and parsing G.I. Bill documents

225

Received 24 March 2023
Revised 7 June 2023
Accepted 13 June 2023

Sara Lafia, David A. Bleckley and J. Trent Alexander
ICPSR, University of Michigan, Ann Arbor, Michigan, USA

Abstract

Purpose – Many libraries and archives maintain collections of research documents, such as administrative records, with paper-based formats that limit the documents' access to in-person use. Digitization transforms paper-based collections into more accessible and analyzable formats. As collections are digitized, there is an opportunity to incorporate deep learning techniques, such as Document Image Analysis (DIA), into workflows to increase the usability of information extracted from archival documents. This paper describes the authors' approach using digital scanning, optical character recognition (OCR) and deep learning to create a digital archive of administrative records related to the mortgage guarantee program of the Servicemen's Readjustment Act of 1944, also known as the G.I. Bill.

Design/methodology/approach – The authors used a collection of 25,744 semi-structured paper-based records from the administration of G.I. Bill Mortgages from 1946 to 1954 to develop a digitization and processing workflow. These records include the name and city of the mortgagor, the amount of the mortgage, the location of the Reconstruction Finance Corporation agent, one or more identification numbers and the name and location of the bank handling the loan. The authors extracted structured information from these scanned historical records in order to create a tabular data file and link them to other authoritative individual-level data sources.

Findings – The authors compared the flexible character accuracy of five OCR methods. The authors then compared the character error rate (CER) of three text extraction approaches (regular expressions, DIA and named entity recognition (NER)). The authors were able to obtain the highest quality structured text output using DIA with the Layout Parser toolkit by post-processing with regular expressions. Through this project, the authors demonstrate how DIA can improve the digitization of administrative records to automatically produce a structured data resource for researchers and the public.

Originality/value – The authors' workflow is readily transferable to other archival digitization projects. Through the use of digital scanning, OCR and DIA processes, the authors created the first digital microdata file of administrative records related to the G.I. Bill mortgage guarantee program available to researchers and the

© Sara Lafia, David A. Bleckley and J. Trent Alexander. Published by Emerald Publishing Limited. This article is published under the Creative Commons Attribution (CC BY 4.0) licence. Anyone may reproduce, distribute, translate and create derivative works of this article (for both commercial and non-commercial purposes), subject to full attribution to the original publication and authors. The full terms of this licence may be seen at <http://creativecommons.org/licenses/by/4.0/legalcode>

This work was supported by a Propelling Original Science (PODS) grant, "Images to Integrated Data", provided by the Michigan Institute for Data Science (MIDAS) at the University of Michigan. The authors are grateful to the Digitization Division staff at the Office of Research Services at the National Archives and Records Administration – and specifically to Denise Henderson – for the guidance and collaboration in the scanning of the paper documents. This research was supported in part through computational resources and services provided by Advanced Research Computing at the University of Michigan, Ann Arbor.

Data availability: Data and code are available on Github, <https://github.com/ICPSR/gi-bill>.



Journal of Documentation
Vol. 79 No. 7, 2023
pp. 225-239
Emerald Publishing Limited
0022-0418
DOI 10.1108/JD-03-2023-0055

general public. These records offer research insights into the lives of veterans who benefited from loans, the impacts on the communities built by the loans and the institutions that implemented them.

Keywords Archives, Digital libraries, Document image processing, Records management, Named entity recognition, Social sciences

Paper type Research paper

Introduction

Digitization has been described as an activity in which information about objects and their context can be converged into a single system (Navarrete and Owen, 2011). Since 2004, digitization has been recognized as a preservation reformatting method by the Association of Research Libraries (ARL) for ensuring continued access to paper-based materials (Arthur *et al.*, 2004). Beyond simply preserving materials, the decision to digitize a collection introduces new possibilities for enhancing access and description of collections. Well-known digitization projects, like Google's campaign to digitize books, have enabled the large-scale analysis of documents by providing new interactions, like zooming in on scanned images and ways to access primary source materials online (Leetaru, 2008).

The digitization of paper-based materials is now standard practice in libraries, archives and museums (Lischer-Katz, 2022). Technical guidelines for digitizing cultural heritage materials describe workflows in which materials are cataloged, scanned, reviewed for quality, archived and published (Federal Agencies Digital Guidelines Initiative, 2022; Puglia *et al.*, 2005). Efforts to embed technologies, like open-source optical character recognition (OCR), into digital historical research workflows (Blanke *et al.*, 2012) have made digitization more relevant, accessible and customizable for adoption by specific research communities. Emerging technologies like deep learning and image segmentation are poised to augment digitization workflows by capturing the structure and content of textual documents (Shen *et al.*, 2021).

Much of the historical documentation OCR literature focuses on the digitization of prose documents or the conversion of hard copy tabular records into digital tabular data (Nagy, 1992; Stančić and Trbušić, 2020). However, bridging these use cases is the digitization of semi-structured historical documents, which hold data that could be converted into a tabular format but are not currently formatted into forms or tables. For example, individual records can be digitized and tabulated for large-scale analysis (Brahney, 2015). Document structures are also critical for maintaining the full meaning of documents, like newspapers and provide valuable context for text mining and historical analysis (Lee *et al.*, 2020). A movement to treat "collections as data" argues that as text is digitized into machine-actionable corpora, including document structure in the text digitization process enables computational research methods such as text mining, data visualization, mapping and network analysis (Padilla *et al.*, 2019).

This paper investigates the feasibility of augmenting a conventional workflow to digitize and parse semi-structured archival records using open-source document image analysis (DIA) and named entity recognition (NER) approaches. We applied these methods to digitize a collection of 25,744 paper-based records from the administration of mortgage guarantees by the Servicemen's Readjustment Act of 1944, commonly known as the G.I. Bill. The output of our process represents the first tabular administrative records dataset available for the study of the implementations and outcomes of the G.I. Bill mortgage guarantee program. We evaluated the performance of multiple OCR methods as well as each text extraction approach by comparing its output against ground truth data. We found that DIA, post-processed using regular expressions, produced the highest quality dataset of structured text. This paper contributes: (1) a digitization workflow for recovering structured text from administrative records; and (2) a novel data collection available to study the G.I. Bill mortgage guarantee program and its beneficiaries.

Background

Text extraction methods

The process of scanning paper-based documents produces representative digital images, which can be indexed for search and retrieval and distributed online. Scanning also enables the conversion of raster images into text through a process known as OCR (Stevens, 1961). OCR was originally developed on typewritten cards to support data entry from paper-based records (Leimer, 1962). Scanned digital images need to support OCR conversion and ensure a high quality of output text (Booth and Gelb, 2006).

Contemporary OCR approaches are more flexible than their predecessors in that they take advantage of document structure to process text blocks (e.g. captions, words in tables) instead of recognizing single characters at a time (Nagy, 1992). Leading OCR engines, like Tesseract, are capable of capturing scanned text with relatively high accuracy, provided that the documents have been correctly prepared and pre-processed (Smith, 2007). OCR engines perform text segmentation and character prediction through classification (Neudecker *et al.*, 2021). Measures such as character accuracy and character error rate (CER) (Neudecker *et al.*, 2021; Rice, 1996) are useful for determining the performance of OCR apart from other steps in digitization workflows. Despite the fact that OCR engines can extract scanned text with high accuracy, the output of OCR is unstructured. In other words, because OCR does not capture the layout context of documents (e.g. columns or fields), separate steps, such as layout analysis, are often needed to detect and evaluate the sources of errors in text extraction (Packer, 2011).

Document image and layout analysis

OCR is just one component of document processing workflows, which also includes page layout analysis and other DIA techniques (Kasturi *et al.*, 2002). Modern DIA methods take advantage of deep learning to classify images and detect document layouts. In recent years, deep learning methods using convolutional neural networks have advanced the state of the art for complex text digitization tasks, such as medieval handwriting classification and layout analysis (Pondenkandath *et al.*, 2017). Strategies and datasets originally developed for computer vision research have been adapted for use in other domains through transfer learning. For instance, the PubLayNet dataset was originally trained to detect the layout of scientific articles and has provided a foundation for developing custom applications for layout detection and analysis of other sources of text (Zhong *et al.*, 2019).

Detecting and incorporating the layout of documents into information extraction tasks makes it possible to retain the layout context of original documents. However, the procedures for adapting and tuning existing document image datasets can be complicated to reproduce due to the use of proprietary services or the need to manage numerous software dependencies. For example, a pipeline was recently developed for digitizing scanned card index records (Amujala *et al.*, 2023). However, its reliance on proprietary OCR and natural language processing (NLP) services available through Amazon Web Services makes it difficult for other researchers to inspect or adapt the underlying models. By contrast, the LayoutParser toolkit supports document image processing and structured text extraction, enabling researchers to adapt existing image layout detection pipelines for custom text extraction tasks (Shen *et al.*, 2021). Techniques like Named Entity Recognition (NER) can further improve the quality and structure of extracted text during post-processing. For example, NER can be used to predict tags for extracted entities (Lu *et al.*, 2013) or recognize semi-structured entities (Irmak and Kraft, 2010).

Materials and methods

Mortgage record index cards

We focus on the digitization of paper-based records from the administration of the mortgage guarantee program of the Servicemen's Readjustment Act of 1944, known as the G.I. Bill, which

guaranteed loans made to U.S. veterans of World War II ([Servicemen's Readjustment Act of 1944, 1944](#)). Between 1944 and 1952, the program guaranteed over two million mortgages ([United States Department of Veterans Affairs, 2013](#)). Prior studies examining the impact of the G.I. Bill have relied on indirect evidence; the literature provides no analysis of administrative records from the implementation of the program ([Katznelson and Mettler, 2008](#)).

The *Index to Loans on Veterans Administration Guaranteed Mortgages, 1946–1954* ("[Index to Loans on Veterans Administration Guaranteed Mortgages, 1946–1954](#)," n.d.) is a collection of 23 linear feet of three-inch by five-inch index cards housed at the National Archives and Records Administration (NARA) in College Park, Maryland. Documents from this collection offer the first administrative data on the execution of the G.I. Bill available to researchers and the general public. Though they are not a comprehensive record of all mortgage guarantee beneficiaries, they constitute a large, novel dataset (n = 25,744 scanned images) that is well-suited for analyzing the long-term impacts of the G.I. Bill program.

Each index card in the collection contains information about the name and address of the mortgagor, the amount of the mortgage, the name of the RFC (Reconstruction Finance Corporation) loan agency approving the mortgage with the issuing agent's serial number and the name and location of the bank handling the mortgage loan ([Zaid, 1973](#)). While most of the index cards contain these common fields, they are hand-typed and the text fields are not in identical positions on each card. [Figure 1](#) illustrates the variation in card layouts. Given that the cards were hand-typed, we expected that OCR methods would convert scanned text to electronically-encoded characters with a high degree of accuracy. The variety of card layouts, however, presented challenges for maintaining the layout context during parsing, using both machine learning and standard approaches.

To evaluate the effectiveness of our text extraction approaches, we developed three truth decks. The first truth deck (n = 100) contained hand-keyed text files for evaluating the quality of the OCR output text. The second truth deck (n = 500) contained cards with regions that we labeled and used as training data for learning card layouts. We labeled batches of 100 cards at a time and trained the model iteratively until we saw only marginal gains in average precision (AP) per label category. [Figure 2](#) shows how we used the open-source software, LabelStudio ([Tkachenko et al., 2020](#)) to label the text fields *Name*, *Location*, *Amount*, *ID*, *Status*, *Agency* and *Other*. These labels provided mappings between text fields and the spatial

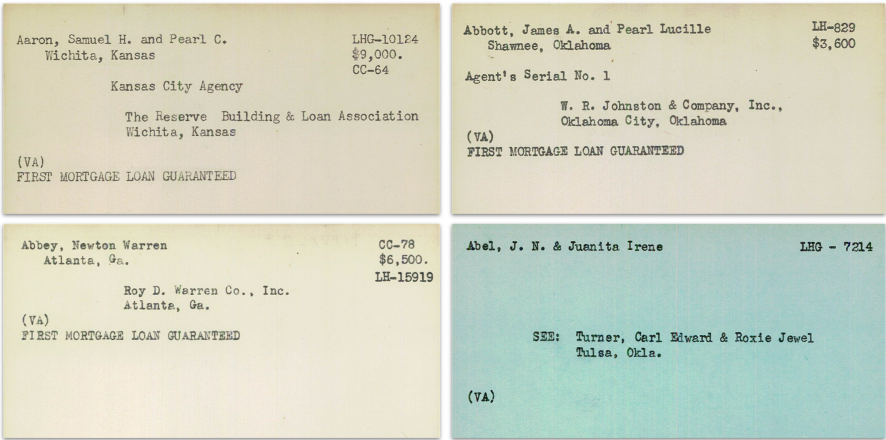


Figure 1.
Example of four
scanned mortgage
record index cards
from the Index to
Loans on Veterans
Administration
Guaranteed Mortgages
(1946–1954)

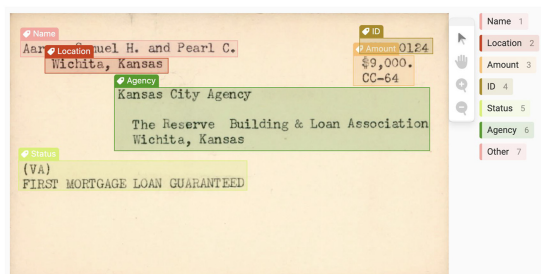
Source(s): Figure by authors

regions of the cards in which they tended to occur. The third truth deck (n = 100) contained hand-keyed data assigned to specific fields to assess the accuracy of parsing.

Digitization and parsing workflow

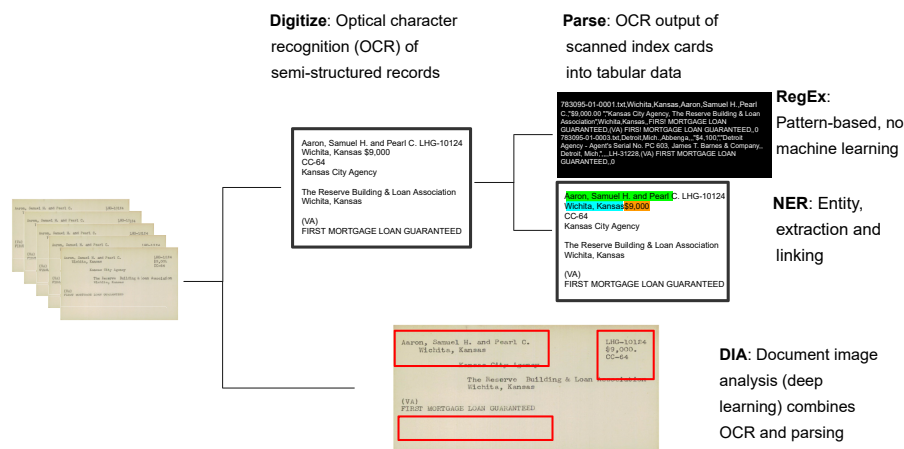
Many archives have card indexes in their repositories. Archival records are often stored as card indexes, which serve as finding aids or contain individual-level information. Historically, the digitization of index cards has been a manual process subject to human error (e.g. inconsistent data entry) and computational error (e.g. inaccurate character recognition) (Amujala *et al.*, 2023). A key challenge of extracting structured textual information from semi-structured historical records is incorporating their layout information into digitization workflows (Shen *et al.*, 2021).

To address this challenge, we employed and assessed multiple methods of digitizing and parsing index card images to develop a workflow that leverages the layout of scanned cards to extract and structure their text. Figure 3 provides an overview of the digitization and parsing methods we used and evaluated in creating our workflow. The workflow transforms the paper-based records into a combined tabular data file, suitable for record linkage and historical analysis. The final model segments each scanned card into regions and predicts a labeled bounding box for each corresponding block of text.



Source(s): Figure by authors

Figure 2. Labeling layout training data indicating text regions with LabelStudio software



Source(s): Figure by authors

Figure 3. Workflow for digitizing and parsing the scanned historical documents

In the first step, *digitization*, the cards were scanned to high-resolution digital images using flatbed scanners by staff at NARA. We then used OCR to extract variably-structured text from the scanned images. To identify the best method for OCR, we tested five standard methods (see Table 1). These included both stand-alone software – (ABBYY FineReader PDF (ABBYY, 2019), Acrobat Pro (Adobe, 2022) and OmniPage Professional (Nuance Communications, Inc., 2011) – and OCR engines called within Python workflows – Tesseract used in LayoutParser (Shen *et al.*, 2021) and Python-tesseract (Hoffstaetter, 2021; Tesseract, 2021). We evaluated the quality of the OCR using the hand-keyed truth deck that we created (see Table 4 in the Results section). In addition to the digital images from NARA, the first step output variably structured text for each card.

Next, we trained a layout model to perform *Document Image Analysis (DIA)* using the LayoutParser library (Shen *et al.*, 2021). We selected Primalayout, a layout analysis model trained on a dataset of magazines and technical and scientific publications, as our base (Antonacopoulos *et al.*, 2009). Our goals for training a custom model were to: (1) recognize distinct blocks of text from index cards as distinct text fields; and (2) classify each text block with its corresponding field label based on its position.

We used the truth deck we created in LabelStudio to train a custom layout detection model. We split the truth deck into 80% training (400 cards) and 20% testing (100 cards) and updated the Primalayout model. To customize our layout detection model with our training data, we used the Fast R-CNN implementation (Girshick, 2015) available in Detectron2, a computer vision library for object detection and image segmentation. We trained each model version using a computing node with a graphics processing unit (GPU) on Great Lakes, a high-performance computing (HPC) platform available for research use through the University of Michigan. It took approximately fourteen hours to process all 25,744 index cards with our final DIA model.

Finally, we compared three approaches for *parsing* the digitized text to an analysis-ready format. The goal of this step was to capture the original structure of the source information (i.e. card layout) in a structured, tabular output file. The three methods for structuring the output that we compared were: (1) regular expressions (*RegEx*) applied to the OCR text output; (2) DIA text bounding box delineation; and (3) NER classification of unstructured text.

First, we crafted RegEx to segment the OCR output. We took advantage of recurring patterns, such as mortgagor names' occurring on the first line and comma delimiting of city/state pairs, to structure the output. Second, we relied on Tesseract, an open-source OCR engine, to extract text within each bounding box predicted by our custom DIA model. This allowed us to associate the label of each bounding box with the extracted text and generate a structured output. Third, we used spaCy, a NLP library (Montani *et al.*, 2020), to predict entity types in the unstructured OCR text output. Using the labels available in the pre-trained English pipeline available in spaCy, we created mappings between the following entity types and fields we defined for the index cards: (1) agency: organization (ORG); (2) amount: money (MONEY); (3) location: countries, cities and states (GPE); (4) name: person (PERSON); and (5)

Table 1.
Optical character
recognition engines
tested

Name	Developer	Version	Operating system
ABBYY FineReader PDF	ABBYY	15	Windows
Acrobat Pro 64-bit	Adobe	2022	Windows
Tesseract in LayoutParser	Open Source	5.2	Linux
OmniPage Professional	Nuance (Kofax)	18	Windows
Tesseract in Python-tesseract	Open Source	5.0	Windows
Source(s): Table by authors			

mortgage id: product (PRODUCT). We report the results from these three parsing approaches in [Table 4](#) in the Results section.

Results

OCR accuracy

We evaluated the output of each of the five OCR methods tested using two measures. To calculate these, we used the PRImA Text Evaluation Tool (OCR Performance Evaluation) 1.5 from Pattern Recognition and Image Analysis (PRImA) Research Lab ([PRImA Research Lab, 2018](#)). First, we looked at character accuracy (the percent of characters not needing to be changed to align with ground truth text), which is one of the most commonly-used measures used to evaluate OCR accuracy, along with its inverse, CER ([Neudecker et al., 2021](#); [Rice, 1996](#)). Second, because of the semi-structured nature of the mortgage records and the relatively low importance of word order from line to line (e.g. it does not matter what order the output text shows the mortgagor's city relative to the mortgage amount), we also measured flex character accuracy ([Clausner et al., 2020](#)), which is well-suited to measure OCR accuracy for materials with complex layouts or with content whose word order is less of a priority than the accuracy within individual words. To average the accuracy across the 100-card sample, we weighted each card based on the number of characters on the card.

[Table 2](#) presents the results of these two measures, both of which range from zero (OCR output differs completely from the truth data) to one (OCR output perfectly matches the truth data). For example, ABBYY FineReader's character accuracy of 0.84 means that 84% of characters do not need to be changed for the OCR to align with the ground truth text, conforming to the word order of the ground truth text. The flex character rate of 0.98 indicates that 98% of characters do not need to be changed, when word order is not a consideration. We found that Tesseract in a pytesseract workflow was the most accurate method of OCR, by both measures, with nearly perfect flex character accuracy. However, ABBYY Finereader and OmniPage both performed very well also, especially using the flex character accuracy measure.

Performance metrics

We also evaluated the performance of our custom layout detection model in the DIA workflow. [Table 3](#) summarizes AP scores for each text field category and for our best model

OCR method	Character accuracy	Flex character accuracy
ABBYY FineReader	0.84	0.98
Adobe Acrobat	0.80	0.84
Tesseract in LayoutParser	0.83	0.87
OmniPage	0.90	0.95
Tesseract in pytesseract	0.95	0.99

Source(s): Table by authors

Table 2.
Character accuracy
and flex character
accuracy for each OCR
method

Agency	Amount	ID	Location	Name	Other	Status	Overall
88.9	71.5	71.0	70.1	66.2	50.2	79.6	71.1

Source(s): Table by authors

Table 3.
Average precision (AP)
per text field category

overall. For example, AP of 88.9 for the text field category of Agency indicates an overlap of 88.9% between ground truth and predicted bounding boxes. Percent overlap is often compared to a threshold value, which is used to distinguish false from true positives. We evaluated model performance based on the overlap (i.e. Intersection over Union, or IoU) between ground truth bounding boxes provided in the held-out test set and bounding boxes predicted by the model. If the overlap exceeded a threshold value (in our case, 50%), the model prediction was considered correct (Padilla *et al.*, 2020). We used this metric to determine how reliably the model's object detection corresponded to that of a human annotator. We trained our model iteratively until we saw only marginal gains in precision (i.e. correct positive predictions) per label category. We were satisfied with our model's ability to identify bounding boxes containing text and correctly label the text region. For example, our final layout model reliably draws a bounding box around the name of a mortgagor on a scanned index card and classifies the text that it contains as a "Name".

Parsing accuracy

After finalizing the custom layout detection model, we compared text parsing performance using regular expressions (Regex), DIA with LayoutParser and named entity recognition (NER) with spaCy. Initially, we only parsed the data into five fields due to limitations in the pre-trained entities in our NER model. We used our hand-keyed truth deck to calculate the CER for each approach [1]. Table 4 summarizes the CER for each method, which indicates the percentage of characters in the parsed output that differ from the truth data, with zero meaning no difference between the OCR output and the truth data (Neudecker *et al.*, 2021). Overall, across all text field categories, we found that text extraction with Tesseract (OCR) and LayoutParser (DIA) was superior to spaCy (NER) and outperformed the use of Regex for several categories, such as "Name" and "Agency".

Both DIA and Regex parsed the OCR data much more accurately than the NER model, and both were able to parse the text into much more granular categories as well. Given that parsing with Regex performed better on some text fields, such as "Location", we also experimented with combinations of approaches. We used Regex to extract cities and states from the "Location" and "Agency" fields. We also split "Name" into first and last names, and if a second person was named, we parsed the name into a separate field. We ultimately compared DIA and Regex parsing on 11 fields.

Table 5 presents the CERs measured for these parsing methods. We found that post-processing DIA text output using Regex resulted in a structured text output with the lowest CER overall. At the field level, each parsing method had strengths and weaknesses, with DIA outperforming Regex in six of 11 fields. The final workflow uses DIA and OCR to extract text from bounding boxes and post-process the output with Regex to parse it into a tabular file with 12 fields (the 11 listed in Table 5 plus the name of the scanned card's image file name).

Table 4.
Character error rate
(CER) for each parsing
approach per text field
category

Field	Regular expressions (Regex)	Document image analysis (DIAs)	Named entity recognition (NER)
Agency	4.41	1.73	56.45
Location	8.98	17.23	35.99
Amount	20.96	33.80	53.66
Name	13.73	5.38	73.37
ID	37.91	15.59	98.50
Overall	10.59	9.49	57.02
Source(s): Table by authors			

Table A1 in Appendix 1 (Supplementary File 1) provides an example of the output file generated from our index digitization and parsing workflow.

Discussion

This paper develops a workflow for digitizing and parsing text from historical paper-based collections, such as index cards. While technologies such as OCR have been available for several decades and make it possible to extract high-quality text from scanned images, off-the-shelf OCR does not yet take advantage of layout information to structure output (Neudecker *et al.*, 2021). Omitting layout contexts from the processing of paper-based collections results in flat, unstructured text. Even if the OCR output text is high quality, it still requires substantial manual processing to delineate separate fields for record linkage and analysis, which are necessary precursors for research use (Stančić and Trbušić, 2020). Manual processes are often the main bottlenecks in digitization workflows (Blanke *et al.*, 2012). Hand transcription, even with the aid of semi-automated tools, limits the volume of data that can be processed and poses the risk of introducing additional sources of human error.

We have incorporated layout detection methods to digitize and parse historical records. The workflow we proposed combines off-the-shelf, state-of-the-art OCR technology with a custom DIA model to create a high-quality, structured text dataset for historical research. Our workflow trains a DIA model without much additional overhead, making it possible for archives to implement it for other index card digitization and related projects. The main requirements for training our DIA model were: (1) the creation of truth decks for validating and training; and (2) access to a HPC environment for training and updating the model. The creation of truth decks was streamlined through the use of existing open-source tools such as LabelStudio. We were also able to access a HPC environment through the University of Michigan. Many academic institutions provide reduced-cost HPC environments for use in academic projects.

Taking advantage of document layout brings computational approaches into closer alignment with human judgments and processes. Computer vision research, in particular, seeks to enable computers to derive information from images and other visual inputs (Zhong *et al.*, 2019). For example, deep learning is already making it possible for models to “learn” the layout of a given document and perform various tasks, such as image segmentation and text extraction (Shen *et al.*, 2021). Incorporating deep learning models into existing records management and digitization efforts in archives holds high potential. The NARA catalog (National Archives and Records Administration, n.d.) has over 6200 index card-based series

Field	RegEx alone	DIA (LayoutParser) with RegEx post-processing
Last name	0.54	4.35
Person 1 name	14.52	6.93
Person 2 name	24.68	4.75
Amount	20.96	33.80
ID	42.14	15.59
City	1.58	24.09
State	16.67	37.06
Agency	4.41	1.73
Agency city	11.53	3.44
Agency state	11.11	10.88
Status	1.95	4.75
Overall	8.84	8.48

Source(s): Table by authors

Table 5.
Character error rate
(CER) for regular
expressions (RegEx)
and document image
analysis (DIA)
approaches per text
field category—more
granular parsing

that may be useful for researchers if they were processed using our digitization (see [Table A2](#) in [Appendix 2](#) (Supplementary File 2) for a few relevant examples. Other archives, libraries and agencies likely have thousands of other similar collections.

The workflow we present brings several advantages to researchers and practitioners. For one, adoption may improve efficiency, freeing up human expertise for other valuable curation tasks, like quality checks and metadata creation, which ensure the discoverability and usability of digital collections. In addition, workflows that leverage deep learning may also support human curators in preparing “collections as data” by making implicit context, like layout information, structurally explicit ([Padilla et al., 2019](#)). For example, in newspaper digitization efforts, the inclusion of layout information and identification of linked entities supports large-scale network analysis and pattern detection ([Lee et al., 2020](#)). Making implicit information (e.g. layout information) explicit increases the analytical utility of the data product for a wider range of scholars and computational research methods.

Conclusion

Paper-based historical records contain rich layout information that must be incorporated to effectively digitize and parse these records into analyzable data. Through the use of deep learning with DIA, we automatically recovered information about document layouts. We applied this technique to process a collection of administrative records related to the Servicemen’s Readjustment Act of 1944, also known as the G.I. Bill. We showed how to use DIA in combination with standard OCR and RegEx approaches to extract high-quality, structured text from scanned images. In summary, this paper contributes: (1) a workflow using scanning, OCR and deep learning to digitize and parse index cards; and (2) a novel, analysis-ready dataset for historical research. The workflow demonstrates the feasibility of incorporating deep learning into archives’ existing digitization and parsing efforts. In addition, the digitized dataset is ready to be linked with additional data sources to further increase its analytical research utility.

This study has produced the first tabular administrative records dataset available to study the implementation and outcomes of the G.I. Bill mortgage guarantee program. While the mortgage program of the Servicemen’s Readjustment Act has been a frequent subject of research across a variety of disciplines, no study to date has used administrative data to examine its impacts. The resulting dataset will improve researchers’ capacity to study the program, its beneficiaries and how this program altered the American way of life. To produce the dataset, we developed and documented a workflow for processing a range of semi-structured archival records. The workflow is applicable to other collections of paper-based or photographic records and uses emerging techniques of DIA based in deep learning. These techniques create new opportunities for the digitization and recovery of historical records—especially records from semi-structured sources.

Note

1. We calculated CER using python code based heavily on xer ([Puigcerver, 2014](#)).

References

- ABBY (2019), “ABBY FineReader PDF (version 15) [computer software]”, available at: https://pdf.abby.com/media/1676/users_guide.pdf
- Adobe (2022), “Acrobat Pro 64-bit (version 2022) [computer software]”, available at: <https://www.adobe.com/acrobat/acrobat-pro.html>

- Amujala, S., Vossmeier, A. and Das, S.R. (2023), "Digitization and data frames for card index records", *Explorations in Economic History*, Vol. 87, 101469, doi: [10.1016/j.eeh.2022.101469](https://doi.org/10.1016/j.eeh.2022.101469).
- Antonacopoulos, A., Bridson, D., Papadopoulos, C. and Pletschacher, S. (2009), "A realistic dataset for performance evaluation of document layout analysis", *2009 10th International Conference on Document Analysis and Recognition*, pp. 296-300, doi: [10.1109/ICDAR.2009.271](https://doi.org/10.1109/ICDAR.2009.271).
- Arthur, K., Byrne, S., Long, E., Montori, C.Q. and Nadler, J. (2004), *Recognizing Digitization as a Preservation Reformatting Method*, Vol. 33 No. 4, pp. 171-180, doi: [10.1515/MFIR.2004.171](https://doi.org/10.1515/MFIR.2004.171).
- Blanke, T., Bryant, M. and Hedges, M. (2012), "Open source optical character recognition for historical research", *Journal of Documentation*, Vol. 68 No. 5, pp. 659-683, doi: [10.1108/00220411211256021](https://doi.org/10.1108/00220411211256021).
- Booth, J.M. and Gelb, J. (2006), *Optimizing OCR Accuracy on Older Documents: A Study of Scan Mode, File Enhancement, and Software Products*, U.S. Government Printing Office, Washington D.C., pp. 1-5, v2.0.
- Brahney, K. (2015), "Information extraction from semi-structured documents MSci. Computer science with industrial experience", available at: <http://miami-nice.co.uk/information-extraction-from-docs.pdf> (accessed 05 June 2015).
- Clausner, C., Pletschacher, S. and Antonacopoulos, A. (2020), "Flexible character accuracy measure for reading-order-independent evaluation", *Pattern Recognition Letters*, Vol. 131, pp. 390-397, doi: [10.1016/j.patrec.2020.02.003](https://doi.org/10.1016/j.patrec.2020.02.003).
- Federal Agencies Digital Guidelines Initiative (2022), *Technical Guidelines for Digitizing Cultural Heritage Materials*, Still Image Working Group, pp. 73-81, No. 3.5, available at: <https://www.digitizationguidelines.gov/guidelines/digitize-technical.html>
- Girshick, R. (2015), "Fast R-CNN", *Proceedings of the IEEE International Conference on Computer Vision*, pp. 1440-1448, available at: http://openaccess.thecvf.com/content_iccv_2015/html/Girshick_Fast_R-CNN_ICCV_2015_paper.html
- Hoffstaetter, S. (2021), "Python-tesseract (version 0.3.8) [computer software]", available at: <https://github.com/madmaze/pytesseract>
- Index to Loans on Veterans Administration Guaranteed Mortgages, 1946 – 1954 (n.d.), "Data set", in *National Archives NextGen Catalog*, available at: <https://catalog.archives.gov/id/783095>
- Irmak, U. and Kraft, R. (2010), "A scalable machine-learning approach for semi-structured named entity recognition", *Proceedings of the 19th International Conference on World Wide Web*, pp. 461-470, doi: [10.1145/1772690.1772738](https://doi.org/10.1145/1772690.1772738).
- Kasturi, R., O'Gorman, L. and Govindaraju, V. (2002), "Document image analysis: a primer", *Sadhana*, Vol. 27 No. 1, pp. 3-22, doi: [10.1007/bf02703309](https://doi.org/10.1007/bf02703309).
- Katznelson, I. and Mettler, S. (2008), "On race and policy history: a dialogue about the G.I. Bill", *Perspectives on Politics*, Vol. 6 No. 3, pp. 519-537, doi: [10.1017/s1537592708081267](https://doi.org/10.1017/s1537592708081267).
- Lee, B.C.G., Mears, J., Jakeway, E., Ferriter, M., Adams, C., Yarasavage, N., Thomas, D., Zwaard, K. and Weld, D.S. (2020), "The newspaper navigator dataset: extracting headlines and visual content from 16 million historic newspaper pages in chronicling America", *Proceedings of the 29th ACM International Conference on Information and Knowledge Management*, pp. 3055-3062, doi: [10.1145/3340531.3412767](https://doi.org/10.1145/3340531.3412767).
- Leetaru, K. (2008), "Mass book digitization: the deeper story of Google books and the open content alliance", *First Monday*, Vol. 13 No. 10, doi: [10.5210/fm.v13i10.2101](https://doi.org/10.5210/fm.v13i10.2101).
- Leimer, J. (1962), "Design factors in the development of an optical character recognition machine", *IRE Transactions on Information Theory*, Vol. 8 No. 2, pp. 167-171, doi: [10.1109/TIT.1962.1057696](https://doi.org/10.1109/TIT.1962.1057696).
- Lischer-Katz, Z. (2022), "The emergence of digital reformatting in the history of preservation knowledge: 1823-2015", *Journal of Documentation*, Vol. 78 No. 6, pp. 1249-1277, doi: [10.1108/JD-04-2021-0080](https://doi.org/10.1108/JD-04-2021-0080).
- Lu, C., Bing, L., Lam, W., Chan, K.I. and Gu, Y. (2013), "Web entity detection for semi-structured text data records with unlabeled data", *International Journal Of. Computational Linguistics and*

- Applications*, Vol. 4 No. 2, pp. 135-150, available at: <http://www.ijcla.org/2013-2/IJCLA-2013-2-pp-135-150-Web.pdf>
- Montani, I., Honnibal, M., Boyd, A., Van Landeghem, S., Peters, H., O'Leary McCann, P., Geovedi, J., O'Regan, J., Samsonov, M., de Kok, D., Orosz, G., Blättermann, M., Altinok, D., Mitsch, R., Kannan, M., Lind Kristiansen, S., Miranda, L., Bournhonesque, R., Baumgartner, P., Hudson, R., Fiedler, L., Daniels, R. and Phatthiyaphaibun, W. (2020), "spaCy: industrial-strength natural language processing in Python (Version v3) [Computer software]", Zenodo, doi: [10.5281/zenodo.1212303](https://doi.org/10.5281/zenodo.1212303).
- Nagy, G. (1992), "At the frontiers of OCR", *Proceedings of the IEEE Institute of Electrical and Electronics Engineers*, Vol. 80 No. 7, pp. 1093-1100, doi: [10.1109/5.156472](https://doi.org/10.1109/5.156472).
- National Archives and Records Administration (n.d), "National archives catalog", available at: <https://catalog.archives.gov/>
- Navarrete, T. and Owen, J.M. (2011), "Museum libraries: how digitization can enhance the value of the museum", *Palabra Clave (La Plata)*, Vol. 1 No. 1, pp. 12-20, available at: <http://www.scielo.org.ar/img/revistas/pacla/v1n1/html/v1n1a03.htm>.
- Neudecker, C., Baierer, K., Gerber, M., Clausner, C., Antonacopoulos, A. and Pletschacher, S. (2021), "A survey of OCR evaluation tools and metrics", *The 6th International Workshop on Historical Document Imaging and Processing*, pp. 13-18, doi: [10.1145/3476887.3476888](https://doi.org/10.1145/3476887.3476888).
- Nuance Communications, Inc (2011), *OmniPage Professional (Version 18) [Computer Software]*, Nuance Communications, Burlington, Massachusetts.
- Packer, T.L. (2011), "Performing information extraction to improve OCR error detection in semi-structured historical documents", *Proceedings of the 2011 Workshop on Historical Document Imaging and Processing*, pp. 67-74, doi: [10.1145/2037342.2037354](https://doi.org/10.1145/2037342.2037354).
- Padilla, T., Allen, L., Frost, H., Potvin, S., Roke, E.R. and Varner, S. (2019), "Always already computational: collections as data: final report", available at: <https://digitalcommons.unl.edu/scholcom/181/>
- Padilla, R., Netto, S.L. and da Silva, E.A.B. (2020), "A survey on performance metrics for object-detection algorithms", *2020 International Conference on Systems, Signals and Image Processing (IWSSIP)*, pp. 237-242, doi: [10.1109/IWSSIP48289.2020.9145130](https://doi.org/10.1109/IWSSIP48289.2020.9145130).
- Pondenkandath, V., Seuret, M., Ingold, R., Afzal, M.Z. and Liwicki, M. (2017), "Exploiting state-of-the-art deep learning methods for document image analysis", *2017 14th IAPR International Conference on Document Analysis and Recognition (ICDAR)*, Vol. 05, pp. 30-35, doi: [10.1109/ICDAR.2017.325](https://doi.org/10.1109/ICDAR.2017.325).
- PRImA Research Lab (2018), "PRImA text evaluation tool (version 1.5) [computer software]", available at: <https://www.primaresearch.org/tools/PerformanceEvaluation>
- Puglia, S.T., Reed, J. and Rhodes, E. (2005), *Technical Guidelines for Digitizing Archival Materials for Electronic Access: Creation of Production Master Files - Raster Images*, National Archives and Records Administration, available at: <https://play.google.com/store/books/details?id=M41NLKdXIdkC>
- Puigcerver, J. (2014), "xer", available at: <https://github.com/jpuigcerver/xer>
- Rice, S.V. (1996), "Measuring the accuracy of page-reading systems", in Nartker, T.A. (Ed.), *University of Nevada, Las Vegas*, available at: <https://www.proquest.com/dissertations-theses/measuring-accuracy-page-reading-systems/docview/304329395/se-2>
- Servicemen's Readjustment Act of 1944 (1944), "78th Congress, Pub. L. 346, 18", available at: [https://hdl-handle-net.proxy.lib.umich.edu/2027/umn.31951d03569283l](https://hdl.handle-net.proxy.lib.umich.edu/2027/umn.31951d03569283l)
- Shen, Z., Zhang, R., Dell, M., Lee, B.C.G., Carlson, J. and Li, W. (2021), "LayoutParser: a unified toolkit for deep learning based document image analysis", *Document Analysis and Recognition – ICDAR*, Vol. 2021, pp. 131-146, doi: [10.1007/978-3-030-86549-8_9](https://doi.org/10.1007/978-3-030-86549-8_9).

-
- Smith, R. (2007), “An overview of the tesseract OCR engine”, *Ninth International Conference on Document Analysis and Recognition (ICDAR 2007)*, Vol. 2, pp. 629-633, doi: [10.1109/ICDAR.2007.4376991](https://doi.org/10.1109/ICDAR.2007.4376991).
- Stančić, H. and Trbušić, Ž. (2020), “Optimisation of archival processes involving digitisation of typewritten documents”, *Aslib Journal of Information Management*, Vol. 72 No. 4, pp. 545-559, doi: [10.1108/AJIM-11-2019-0326](https://doi.org/10.1108/AJIM-11-2019-0326).
- Stevens, M.E. (1961), *Automatic Character Recognition: A State-Of-The-Art Report*, U.S. Department of Commerce, available at: <https://hdl-handle-net.proxy.lib.umich.edu/2027/mdp.39015077289836>
- Tesseract (2021), “Tesseract OCR (version 5.0) [computer software]”, available at: <https://github.com/tesseract-ocr/tesseract>
- Tkachenko, M., Malyuk, M., Shevchenko, N., Holmanyuk, A. and Liubimov, N. (2020), “LabelStudio: Data labeling software (version 1.7) [computer software]”, available at: <https://github.com/heartexlabs/label-studio>
- United States Department of Veterans Affairs (2013), “History and timeline—education and training”, available at: <https://www.va.gov/education/about-gi-bill-benefits/>
- Zaid, C. (1973), *Preliminary Inventory of the Records of the Reconstruction Finance Corporation, 1932-1964*, National Archives & Records Service, (Record Group 234), available at: <https://hdl-handle-net.proxy.lib.umich.edu/2027/uiug.30112101560024>
- Zhong, X., Tang, J. and Jimeno Yepes, A. (2019), “PubLayNet: largest dataset ever for document layout analysis”, *2019 International Conference on Document Analysis and Recognition (ICDAR)*, pp. 1015-1022, doi: [10.1109/ICDAR.2019.00166](https://doi.org/10.1109/ICDAR.2019.00166).

(The Appendix follows overleaf)

Table A1.
Example of DIA layout
model text output post-
processed with regular
expressions

Appendix 1
(Supplementary file 1): example output

Image	City	State	Last name	Person 1 name	Person 2 name	Amount	Agency	Agency city	Agency state	ID	Status
783095-01-0003.jpg	Detroit	Mich	Abbenga	Arnold N	Geraldine	\$4,100	Detroit Agency - Agent's Serial No. PC 603, James T, Barnes and Company	Detroit	Mich	LH- 31228	(V4) FIRST MORTGAGE LOAN GUARANTEED
783095-01-0004.jpg	Tulsa	Oklahoma	Abbey	Leonard Ray	Barbara Joan	\$8,000	Agent's Serial No. 44, W. R. Johnston and Co., Inc	Oklahoma City	Oklahoma	LH- 5746	(V4) FIRST MORTGAGE LOAN GUARANTEED
Source(s): Table by authors											

Appendix 2

(Supplementary file 2): index card collections

Digitizing and
parsing G.I. Bill
documents

239

Title	NARA ID	Description
Card Index of Licensees, ca 1918–ca 1918	5111291	This series consists of a card index listing businesses (sometimes listed under the name of an individual, otherwise listed by the name of the firm) licensed by the Iowa State Food Administration. Both wholesalers and retailers are included. The name and address of the establishment is listed, along with an unidentified sequence of numbers, the last of which is the number of the form on which the business was obliged to report to the United States Food Administration
Awards Files Card Index, 1944– 1945	611132	This series consists of an index of the awards and decorations given to soldiers of the 44th Infantry Division during the campaigns in northern France, the Rhineland and central Europe during World War II. The awards referenced in this series include the Silver Star, Bronze Star, Air Medal and Purple Heart. Each index entry provides the name of the soldier given the award, the type of award given and the date of the award
Publications Card Files, 1944– 1945	624400	This series consists of cards showing the name of a propaganda publication or leaflet; the date of publication; the publisher; the number of copies printed; the date of pickup for distribution; and the date, number of copies and location of dissemination. Many of the cards have a copy of the reference publication attached
Card Index of Trusts, 1917– 1934	6879969	This series consists of an index for trusts established for individuals or companies whose property was seized by the Alien Property Custodian. The cards in this series include spaces for name and address, trust number, date opened, report number, ticket number and other information. Sometimes, only a name and trust number are included on the card
Death Certificate Card Index, 1914–February 1915	7408557	The series contains an index that records deaths that occurred in the Canal Zone. The records include information concerning the deceased such as name, age, color, sex, nationality, occupation and employment, residence, address, nature of illness and cause of death, attending physician, date of death and grave number

Source(s): Table by authors, [National Archives and Records Administration \(n.d.\)](#)

Table A2.
Examples of index card
collections listed in the
National Archives and
Records
Administration
(NARA) Catalog—all
information copied
directly from NARA
Catalog website

Corresponding author

Sara Lafia can be contacted at: slafia@umich.edu

For instructions on how to order reprints of this article, please visit our website:

www.emeraldgroupublishing.com/licensing/reprints.htm

Or contact us for further details: permissions@emeraldinsight.com