

Inhaltliche Erschließung

Ramona Bräu*, Kerstin Hofmann, Sonja Nilson und Christa Zwilling-Seidenstücker

#EveryNameCounts – Die Crowdsourcing-Initiative der Arolsen Archives

Ein Erfahrungsbericht

<https://doi.org/10.1515/iwp-2021-2157>

„As I type each card into the database I think of each name and the person it belonged to – and I honor them as a human being.“¹

Zusammenfassung: Das Crowdsourcing-Projekt #EveryNameCounts der Arolsen Archives begann 2020 als Schulkampagne und wuchs binnen weniger Monate zu einer weltweiten Initiative. Durch die Arbeit der Freiwilligen werden Millionen von Datensätzen zu Konzentrationslagerdokumenten erfasst. Diese Datengrundlage optimiert einerseits die Suche im Online-Archiv der Arolsen Archives und macht andererseits eine auf Massendaten gestützte Forschung erst möglich. Die Arbeit der Freiwilligen geht dabei über reine Datenerhebung weit hinaus. Es entsteht ein digitales Denkmal.

Deskriptoren: Archiv, Inhaltliche Erschließung, Online, Digitalisierung, Projekt, Geschichte, Zusammenarbeit, Crowdsourcing

#EveryNameCounts – The crowdsourcing initiative of the Arolsen Archives. A field report

Abstract: The Arolsen Archives' crowdsourcing project #EveryNameCounts began in 2020 as a school campaign and within a few months grew into a worldwide initiative. Through the work of volunteers, millions of records of concentration camp documents are captured. On the one

hand, this data basis optimizes the search in the online archive of the Arolsen Archives and, on the other hand, makes research based on mass data possible in the first place. The work of the volunteers goes far beyond mere data collection. A digital monument is being created.

Descriptors: Archive, Indexing, Online, Digitisation, Project, History, Cooperation, Crowdsourcing

#EveryNameCounts – L'initiative de crowdsourcing des Arolsen Archives. Un rapport de terrain

Résumé: Le projet de crowdsourcing des Arolsen Archives, #EveryNameCounts, a commencé comme une campagne scolaire en 2020 et s'est transformé en une initiative mondiale en quelques mois. Grâce au travail bénévoles, des millions des données de documents de camps de concentration sont saisies. D'une part, cette base de données optimise la recherche dans les archives en ligne des Arolsen Archives et, d'autre part, rend possible en premier lieu la recherche basée sur des données de masse. Le travail des volontaires va bien au-delà de la simple collecte de données. Un monument numérique est en cours de création.

Descripteurs: Archive, Indexation du contenu, Numérisation, Projet, Histoire, Coopération, Crowdsourcing

¹ <https://www.zooniverse.org/projects/arolsen-archives/every-name-counts/talk/3279/1464353?comment=2376649&page=1>

***Kontaktperson:** Ramona Bräu, Arolsen Archives, Große Allee 5–9, 34454 Bad Arolsen, E-Mail: ramona.braeu@arolsen-archives.org
Kerstin Hofmann, Arolsen Archives, Große Allee 5–9, 34454 Bad Arolsen, E-Mail: kerstin.hofmann@arolsen-archives.org
Sonja Nilson, Arolsen Archives, Große Allee 5–9, 34454 Bad Arolsen, E-Mail: sonja.nilson@arolsen-archives.org
Christa Zwilling-Seidenstücker, Arolsen Archives, Große Allee 5–9, 34454 Bad Arolsen, E-Mail: christa.seidenstuecker@arolsen-archives.org

Die Indexierung von Massendaten, die Anreicherung von vorhandenen Metadaten und die Aufbereitung und Verknüpfung von Datensätzen gehören zu den umfangreichsten wie aufwendigsten Herausforderungen der modernen Archivarbeit. Die Anforderungen an die Benutzbarkeit von Archiven steigt entsprechend den Gewohnheiten derjenigen, die sie in einer zunehmend digitalisierten Alltagswelt nutzen. Digitalisierung endet aber – entgegen einer weit verbreiteten Vorstellung – nicht mit dem Scannen von Dokumenten, sondern beginnt eigentlich erst mit der Erstellung des Digitalisats. Crowdsourcing-Projekte werden in diesem Zusammenhang häufig als reines Mittel zum

Zweck gesehen, um die Datengrundlage für diverse Such- und Filterfunktionen zu erheben, die nicht einfach durch die Integration von Normdaten optimiert werden können. Sie versprechen geringe Kosten, schnelle Bearbeitung und Ergebnisse, die nicht auf die Fähigkeiten eines Algorithmus begrenzt sind. Allerdings verkennt dieser doch sehr pragmatische, ja rationale Ansatz die eigentlichen Optionen und Chancen des Crowdsourcings sowie die Anforderungen an die Projektbeteiligten. Die Teilhabe an und das Engagement für ein solches Projekt hängen für die Freiwilligen der Crowd zu einem großen Teil von der Sinnhaftigkeit und von einem erkennbaren individuellen oder gesellschaftlichen bzw. wissenschaftlichen Mehrwert ab. Freiwillige ermüden in der Praxis schnell, wenn die Usability zu wünschen lässt. Dieser emotionale Charakter eröffnet aber zugleich einen Raum der Interaktion und macht ein solches Projekt über die reine Datenerfassung hinaus zu einem Seismometer der Relevanz.

Durch die historisch bedingte, institutionelle Verfasstheit der Arolsen Archives und die Tatsache, dass bereits ein Großteil der Dokumente in digitalisierter Form vorliegen, waren die Grundvoraussetzungen für ein solches Crowdsourcing-Projekt von vornherein günstig. Die Arolsen Archives, vormals International Tracing Service (ITS), sind ein internationales Zentrum über NS-Verfolgung mit dem weltweit umfassendsten Archiv zu den Opfern und Überlebenden des Nationalsozialismus. Neben den 30 Millionen Originaldokumenten und weiteren Millionen von Dokumentenkopien umfasst die stetig wachsende Datenbank insgesamt 110 Millionen Objekte. Die Sammlung mit Hinweisen zu rund 17,5 Millionen Menschen gehört zum UNESCO-Weltdokumentenerbe. Sie beinhaltet Dokumente zu den verschiedenen Verfolgten Gruppen des NS-Regimes, zur Zwangsarbeit sowie zu Displaced Persons (DPs) und Migration nach 1945. Die Arolsen Archives haben eine digitale Strategie 2021 bis 2025 entwickelt, um den weltweiten barrierefreien Zugriff auf den gesamten Dokumentenbestand zu ermöglichen. Derzeitig sind im Online-Archiv bereits über 26 Millionen Dokumente einsehbar.²

Die Arolsen Archives begreifen sich zudem als Informations-Hub und möchten gemeinsam in einem Netzwerk, das sowohl die Copyholder³ als auch Gedenkstätten, Erinnerungsorte und Museen, aber explizit auch die uni-

versitäre Forschung sowie das Engagement kleinerer Einrichtungen (Vereine, Initiativen und individuelle Multiplikatoren) einbezieht, die Möglichkeit schaffen, Sammlungen und Mikroarchive zu digitalisieren und online zugänglich zu machen. Neben der digitalen Sicherung von Informationen und Dokumenten für die Zukunft stehen die archivische Indexierung, Erschließung und Kontextualisierung solcher Überlieferungen im Vordergrund, die wiederum Teil sowie notwendige Basis für Projekte im Bereich der Data Science, der Digital Humanities bzw. der Public Science und Citizen Science sind.

Die Erfassung der Daten im Rahmen der Initiative #EveryNameCounts erfolgt über die international größte nichtkommerzielle Citizen Science-Plattform Zooniverse (mit Sitz an der Oxford University und dem Adler-Planetarium in Chicago, Illinois). Die Plattform ist frei zugänglich und hunderttausende Freiwillige aus der ganzen Welt helfen mit, große Datenbestände für Wissenschafts- und Kultureinrichtungen zu erschließen. Bei einem Vorhaben wie #EveryNameCounts auf einer solchen reichweitenstarken Plattform ist die Frage nach den Veränderungen und Entwicklungen in der Gedenk- und Erinnerungskultur am Ende der Zeitzeugenschaft ebenso evident, wie auch das Hinterfragen der Möglichkeiten und Grenzen eines Gedenkens mit digitalen Mitteln. Für die Überlebenden und Angehörigen der Opfer hat sich die Bedeutung der Sammlungen der Arolsen Archives über die Jahre hin zu einem „Denkmal aus Papier“⁴ entwickelt. Die Arbeit der Freiwilligen mit den Dokumenten auf Zooniverse ergänzt diese Wahrnehmung um eine weitere bedeutsame Facette: das Digitale.

„I have done a lot of genealogical transcription in the past and somehow didn't realize there are other ways to do this kind of work to help connect people with their history and their ancestors. I'm very excited to have found this platform to participate in new ways. I have always loved history and feel strongly about the need to preserve and make available this kind of information. I'm looking forward to contributing to this project and many others. Thanks for having me!“⁵

Im Folgenden werden die Entwicklungsstufen und Erfahrungen des Projektes chronologisch nachgezeichnet und besondere Herausforderungen an den Projekthinhaber Arolsen Archives beschrieben.

² Vgl. <https://collections.arolsen-archives.org/>.

³ Die Partnerinstitutionen der Arolsen Archives sind: Archives de l'État en Belgique, Brüssel (Belgien), Archives Nationales, Pierrefitte-sur-Seine (Frankreich), Centre de Documentation et de Recherche sur la Résistance (Luxemburg), Instytut Pamięci Narodowej (IPN), Warschau (Polen), The Wiener Library, London (Vereinigtes Königreich),

Yad Vashem, Jerusalem (Israel), US Holocaust Memorial Museum (USHMM), Washington (USA) und das Bundesarchiv (Deutschland)

⁴ Vgl. Borggräfe, Henning, Höschler, Christian, Panek, Isabel (Hrsg.): Ein Denkmal aus Papier. Die Geschichte der Arolsen Archives, Bad Arolsen 2019.

⁵ <https://www.zooniverse.org/projects/arolsen-archives/every-name-counts/talk/3279/1557243?comment=2521998&page=1>

Von der Schulkampagne zur weltweiten Initiative

#EveryNameCounts startete am 27. Januar 2020 als lokales Bildungsprojekt. Unter dem Motto „jeder Name zählt...!“ erfassten rund 1.000 Schülerinnen und Schüler aus Hessen einen Tag lang Namen und Geburtsdaten von Deportations- und Transportlisten. Um die Arbeit mit den Dokumenten aus der NS-Zeit vor- und nachzubereiten, erhielten die teilnehmenden Schulen umfassendes Unterrichtsmaterial. Das Feedback war durchweg positiv: Die Konfrontation mit den Originalquellen mit Namen, Alter und Wohnorten der Opfer war für die Schülerinnen und Schüler ein ungewohnt direkter Zugang zu den Verbrechen des NS-Regimes. Sie setzten sich durch das Projekt intensiv mit einzelnen Schicksalen auseinander und leisteten zugleich einen greifbaren Beitrag gegen das Vergessen. Die Menge und Qualität der Metadaten, die so an nur einem Tag erfasst worden waren, machte außerdem deutlich, dass das Projekt nicht nur pädagogischen Wert hat. Crowdsourcing zeigte das Potential, in Zukunft ein entscheidender Faktor bei der Indexierung der Dokumente der Arolsen Archives zu sein.

Als die Corona-Pandemie ab dem Frühjahr 2020 weite Teile des Lebens in den digitalen Raum verlegte, wurde dieser Prozess beschleunigt: Viele Gedenkveranstaltungen zum 75. Jahrestag der Befreiung ehemaliger Konzentrationslager konnten nicht vor Ort stattfinden. Dies gab den Anstoß, das Projekt kurzfristig auszubauen und so weltweit als einen digitalen Ort des Gedenkens zugänglich zu machen. Aus „jeder Name zählt“ wurde #EveryNameCounts. Die positiven Reaktionen waren sowohl in den klassischen Medien als auch in den Social Networks bemerkenswert. Die Bereitschaft zur Beteiligung war enorm:

“Hello, I am Mia from the US. I have only begun working on this project a few days ago, after having read about it in the New York Times.”⁶

Heute ist #EveryNameCounts eine internationale Initiative, die aktives Gedenken an die Opfer der nationalsozia-

listischen Verfolgung mit historischem Lernen verbindet. Über 19.000 registrierte Freiwillige arbeiten am Bau des „digitalen Denkmals“ mit.

“As I document each card, at first I merely transcribed the data needed for the archives. Then as days have passed, I started re-searching the camp admittance form. The cards reveal so much more than just the names, date of birth and prisoner number. Often prisoners were transferred to other camps and that is noted or sadly their date of death appears. [...] Categories of nationalities can also be found. As each name appears on my screen, a snapshot of who they were becomes apparent – making them even more precious and missed.”⁷

Die Arbeit mit der „Crowd“ bringt allerdings auch Herausforderungen mit sich: Bei den Freiwilligen kann kein Hintergrundwissen vorausgesetzt werden und nur ein Teil verfügt über Kenntnisse der deutschen Sprache, in der die Dokumente verfasst sind. Aus diesem Grund eignen sich besonders gleichförmige Dokumente für #EveryNameCounts. Die Freiwilligen können aus unterschiedlichen „Workflows“ auswählen. Diese unterscheiden sich in der Art der Daten, die erfasst werden und der Herkunft der Dokumente. Aktuell werden vor allem individuelle Unterlagen, wie Häftlings-Personal-Karten und Häftlingspersonalbögen aus verschiedenen Konzentrationslagern, bearbeitet. Konkret werden Daten wie Name, Geburtsdatum, Geburtsort, Beruf, Staatsangehörigkeit, Familienstand, Religion, letzter Wohnort, Häftlingsnummer oder Haftkategorie erfasst. Aber auch Angaben zu Angehörigen, die oft ebenfalls Verfolgte waren, werden eingetragen oder Informationen zu Überstellungen, mit deren Hilfe Verfolgungswege rekonstruiert werden können.

Die Dokumente werden einzeln angezeigt und die zugehörigen Daten in eine Maske eingegeben. In die Felder kann entweder Freitext eingetragen oder die zutreffenden Angaben aus einem Dropdown-Menü ausgewählt werden. Zu jedem Feld gibt es eine ausführliche bebilderte Anleitung, wo die jeweilige Information im Dokument zu finden ist und wie sie korrekt erfasst werden soll. Außerdem steht Freiwilligen seit dem 27. Januar 2021 eine digitale Einführung zur Arbeit im Projekt zur Verfügung.⁸ Um Fehler möglichst ausschließen zu können, werden die Daten zu jedem Dokument mindestens dreimal erfasst. Anschließend werden die Daten ausgewertet und einer Qualitätssicherung unterzogen.

Im Forum haben Freiwillige die Möglichkeit, Fragen zu stellen und Rückmeldungen zum Projekt zu geben. Die-

⁶ Vgl. Curry, Andrew: How Crowdsourcing Aided a Push to Preserve the Histories of Nazi Victims, in: New York Times (Online-Ausgabe) v. 3.6.2020, abrufbar: <https://www.nytimes.com/2020/06/03/world/europe/nazis-arolsen-archive.html>; Ders.: Found in the Crowd: Preserving the Histories of the Nazis' Victims, in New York Times (Print-Ausgabe), Section A, 4.6.2020, S. 11. Zitat: <https://www.zooniverse.org/projects/arolsen-archives/every-name-counts/talk/3279/1457620?comment=2366054&page=1>

⁷ <https://www.zooniverse.org/projects/arolsen-archives/every-name-counts/talk/3279/1464353?comment=2402103&page=1>

⁸ Vgl. <https://arolsen-archives.org/enc-intro/de/>.

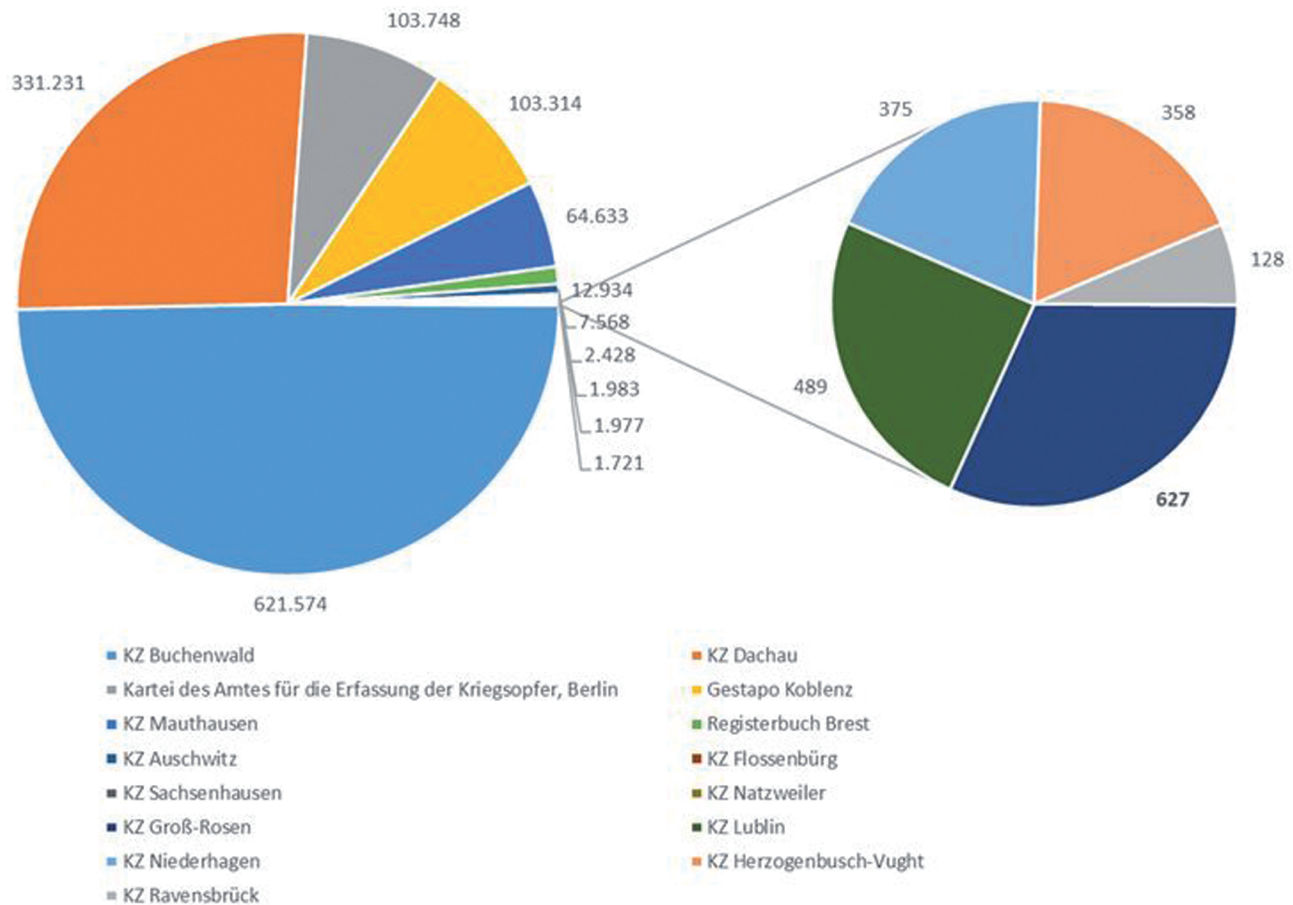


Abb. 1: Anzahl der Dokumente nach Beständen.

ses Feedback fließt wieder in die Konstruktion der neuen Workflows ein. Da wir im Forum beispielsweise immer wieder Korrekturen von fehlerhaften Ortsangaben in den Dokumenten erhielten, entschlossen wir uns, einen speziellen Workflow zur Erfassung von GEO-Daten anzubieten. Darin haben Freiwillige die Möglichkeit, die Daten auf eine Weise mit uns zu teilen, die uns eine gezielte Weiterverarbeitung ermöglicht.

Zum 27. Januar 2021 wurden weitere 600.000 individuelle KZ-Dokumente über #EveryNameCounts bereitgestellt. Damit beläuft sich die Gesamtzahl auf etwa 1,6 Millionen Dokumente aus 15 Beständen, die in 31 Workflows bearbeitet wurden bzw. werden. Mit der Indexierung per Crowdsourcing in dieser Größenordnung haben die Arolsen Archives Neuland betreten. Seit dem Projektstart vor über einem Jahr haben wir viel dazu gelernt und das Projekt hat sich beständig weiterentwickelt.

Bereitstellung der Daten

Um die Arbeit in den verschiedenen Workflows überhaupt zu ermöglichen, ist zunächst die Art und Weise der Bereitstellung von Daten auf der Plattform ausschlaggebend. Begonnen haben wir mit den Mitteln, die Zooniverse bereitstellt: Der Zooniverse Project Builder ist eine graphische Oberfläche zur Gestaltung der Projektseite. Darüber werden u. a. auch die zu bearbeitenden Bilder hochgeladen.

In den ersten Monaten wurden die Bilddateien direkt vom Server hochgeladen, was durch den Erfolg des Projektes bald an Grenzen stieß. Zooniverse empfiehlt für den Upload, Bilddateien in Gruppen von 500 bis 1.000 Stück zu bündeln. Bei Zehntausenden von Bildern oder mehr ist diese Vorgehensweise nicht nur zeitaufwändig, sondern auch fehleranfällig. Deshalb war es folgerichtig, dass wir uns technisch neu aufgestellt haben. Seit November 2020 werden die Bilddateien nicht mehr vom Server, sondern direkt aus der institutionseigenen Cloud nach Zooniverse hochgeladen.

Ein Nebeneffekt dieses Verfahrens kam ebenfalls zum Tragen: Die heimischen DSL-Anschlüsse konnten entlastet

Untitled subject set #93978

A subject is a unit of data to be analyzed. A subject can include one or more images that will be analyzed at the same time by volunteers. A subject set consists of a list of subjects (the “manifest”) defining their properties, and the images themselves.

Feel free to group subjects into sets in the way that is most useful for your research. Many projects will find it's best to just have all their subjects in 1 set, but not all.

The project has 1040151 uploaded subjects. You have uploaded 1590000 subjects from an allowance of 2000000. Your uploaded subject count is the tally of all subjects (including those deleted) that your account has uploaded through the project builder or Zooniverse API. Please [contact us](#) to request changes to your allowance.

We strongly recommend uploading subjects in batches of 500 - 1,000 at a time. When uploading large numbers of subjects, we recommend using our [Panoptes command line interface](#) or our [Panoptes Client package for Python](#) rather than the web portal.

NAME

Untitled subject set

A subject set's name is only seen by the research team.

This set contains 0 subjects:

Page of ?

Drag-and-drop or click to upload manifests and subject images here (you must select the media files as well as the manifest)

Manifests must be .CSV or .TSV. The first row should define metadata headers. All other rows should include at least one reference to an image filename in the same directory as the manifest.

Headers that begin with “#” or “/” denote private fields that will not be visible to classifiers in the main classification interface or in the Talk discussion tool.

Headers that begin with “!” denote fields that **will not** be visible to classifiers in the main classification interface but **will be** visible after classification in the Talk discussion tool.

Subject images can be up to 1000KB and any of: .jpg, .jpeg, .png, .gif, .svg, .mp3, .m4a, .mpeg, .txt, .json and may not contain /, \, ., : , ,

Upload 0 new subjects

Abb. 2: Uploadfunktion des Zooniverse Project Builders.

werden, denn alle Beschäftigten der Arolsen Archives arbeiten seit März 2020 im Homeoffice. Diese Umstellung führte allerdings dazu, dass der Project Builder nicht mehr genutzt werden konnte. Zooniverse bietet als Lösung zwei andere Wege an, um große Mengen von Dateien hochzuladen: Über das Kommandozeilenprogramm „Panoptes CLI“⁹ oder den auf Python beruhenden „Panoptes Client“¹⁰, wobei wir uns für die Nutzung der Kommandozeile entschieden haben. Seitdem können Massenaufloads einer theoretisch unbegrenzten Menge von Bilddateien durchgeführt werden.

In der Praxis hat es sich als sinnvoll erwiesen, nicht mehr als 40.000 Dokumente in einem Paket hochzuladen.

Neben der Entscheidung, jedem Digitalisat Metadaten beizufügen, wurde auch der Workflow bezüglich der Bilder angepasst. Denn die Erfahrung der ersten Monate hat gezeigt, dass eine Menge homogener und gut lesbarer Bilder in einem Subject Set, d.h. einer definierten Gruppe von Bildern, die Arbeit der Freiwilligen erleichtert, weniger Kommentare im Forum hervorruft und die Qualität der eingegebenen Daten deutlich verbessert. Deshalb wurden fortan anstatt ganzer Bestände mehrheitlich zuvor per Clustering gebildete Gruppen von Dokumenten für den Upload gebündelt. So haben wir beispielsweise aus dem Bestand „Konzentrationslager Buchenwald“ den Teilbestand mit individuellen Unterlagen¹¹ ausgewählt und darin alle Häftlings-Personal-Karten sowie Häftlingsper-

⁹ <https://github.com/zooniverse/panoptes-cli>

¹⁰ <https://github.com/zooniverse/panoptes-python-client>

¹¹ <https://collections.arolsen-archives.org/archive/1-1-5-3/?p=1>

SUBJECT METADATA

File ID	6726800
Collection	1.1.5.3 - Individual Documents male Buchenwald
Document ID	6726802
Database entry	Zygmunt NOWAK, born 03.06.1917
Online archive	https://collections.arolsen-archives.org/en/archive/6726800/?p=1&s=6726800&doc_id=6726802

Abb. 3: Anzeige „Subject Metadata“ als Resultat eines Manifestes.

sonalbögen geclustert, exportiert und für diese zwei Arten von Dokumenten jeweils einen Workflow angelegt. Da aus früheren Projekten gute Erfahrungen mit der Layout-Erkennung von Dokumenten vorlagen und die notwendigen Prozesse bereits aufgesetzt waren, erfolgte die Filterung passender Dokumente ohne Zeitverzug.

Ein weiterer wichtiger Schritt war die Einbindung der auf dem Server liegenden Digitalisate in eine Qualitätssicherung noch vor dem Upload. Es hatte sich nämlich als durchaus störend erwiesen, wenn den Freiwilligen immer wieder Bilder zur Indexierung angezeigt wurden, die aus unterschiedlichen Gründen nicht passten¹²: Teilweise waren Seiten leer oder schwarz. Ebenso konnten die Deck- oder Zwischenblätter, deren Erfassung im Rahmen der bestehenden Workflows nicht zielführend war, herausgefiltert werden. Bei den Dokumenten, die mit Vorder- und Rückseite als ein Zooniverse-Subject hochgeladen worden waren, kam es außerdem vor, dass es in Folge einer Seitenverschiebung, bei der die Vorderseite fehlte, die Rückseite vorrückte und eine weitere Leerseite als Rückseite erschien.

Dieses durchaus häufige Problem geht allerdings auf einen fehlerhaften Import innerhalb unseres Archivinformationssystems zurück. Bei einem Umfang von über 100 Millionen digitalen Objekten können bei automatisierten Prozessen nicht alle Fehler aufgefangen werden. Ein weiterer Vorteil, den das Zooniverse-Projekt mit sich bringt, ist also die Möglichkeit, solche Auffälligkeiten gezielt nachzuarbeiten.

Neben der Optimierung der Qualität der bereitgestellten Digitalisate wurde außerdem die Verfahrensweise bzgl. des Umfangs der beigestellten Metadaten weiterentwickelt. Zooniverse bietet die Möglichkeit, beim Upload von Bilddateien diesen in Form eines Manifests auch Metadaten beizufügen. Eine CSV-Datei ist alles, was dafür

nötig ist. Im Herbst 2020 sind wir dazu übergegangen, die Dokumente durchgängig mit Metadaten zu versehen. Da der Upload über „Panoptes CLI“ in jeden Fall ein Manifest erfordert, war die Bereitstellung spätestens ab November 2020 auch technisch zwingend. Darüber hinaus ergab sich im Besonderen mit Blick auf die Nutzungserfahrungen aus der historischen Kontextualisierung und der Verortung der Dokumente innerhalb der Archivtektonik ein klarer Mehrwert. Folglich werden seither jedem Dokument die Dokumenten-ID¹³ sowie Nummer und Name des Bestandes als weiterführende Informationen beigelegt. Wo vorhanden, werden ebenfalls Personendaten aus der institutionseigenen Datenbank hinzugefügt. Hierfür wird per SQL-Abfrage des Archivinformationssystems für den in Bearbeitung befindlichen Bestand eine Werteliste der Felder Vorname, Nachname und Geburtsdatum sowie gegebenenfalls Häftlingsnummer ausgegeben.

Wenn die Dokumente bereits in unserem Onlinearchiv¹⁴ veröffentlicht wurden und es möglich ist, eine direkt auf das Dokument verweisende URL zu erzeugen¹⁵, wird dem Manifest zusätzlich ein Hyperlink hinzugefügt. Die Liste der Dateinamen wird mit Hilfe der Software „Total Commander“ vom Server ausgelesen. Beide Dateien werden als Projekte im Tool „OpenRefine“ angelegt und die Daten über ein Matching der Dokumenten-ID zusammengeführt.

Community-Management

„It feels good to be helpful in a concrete way and work as part of a team, in a way giving voice to, and standing up for those who suf-

¹² Im Bereich „Diskutieren“ können Freiwillige über das Setzen des Hashtags #skipped unser Team gezielt auf solche Dokumente hinweisen, die wir dann per „Panoptes CLI“ künstlich auf den Status „fertig bearbeitet“ setzen können.

¹³ Wenn die Digitalisate nachgenutzt und deshalb zitiert werden, dient die Dokumenten-ID als eindeutiger Identifikator der Archivalie.

¹⁴ <https://collections.arolsen-archives.org/>

¹⁵ Wenn es zu einem Dokument noch keine Personendaten als Metadaten gibt, dann ist eine direkte Adressierung über eine eindeutige URL häufig nicht möglich.

*ferred under the National Socialist regime, and perhaps making it easier for a family member, or friend even, to find them.*¹⁶

Eine nicht zu unterschätzende Kernaufgabe der Projektleitung eines Crowdsourcing-Vorhabens ist es, die wachsende Community zu betreuen und die Ressourcen vorab bei der Projektplanung zu bedenken. Die Erwartungshaltung der Freiwilligen an den Support – sowohl den technischen als auch den inhaltlichen – ist ebenso zentral für das Gelingen des Projekts wie eine durchdachte Nutzerführung mit einem getesteten Usability-Konzept in den Workflows. Anfänglich genügte eine Kollegin, die sich vollumfänglich mit #EveryNameCounts beschäftigte. Es zeigte sich jedoch bereits nach wenigen Monaten, dass das Projekt und vor allem die Betreuung der Freiwilligen und ihrer Rückfragen nicht mehr allein von einer Person zu leisten war.

Während es im April 2020 noch 38.624 Klassifikationen und 4.445 Forumsbeiträge waren, stieg die Anzahl der eingegebenen Klassifikationen im Juni 2020 auf 746.548 und 12.765 Beiträge im Forum.¹⁷ Im Juli 2020 wurde deshalb ein Moderationsteam aus sechs Personen gegründet. Die internen Arbeitsabläufe mussten sich erst einpendeln, doch schnell zeigte es sich, dass dies am besten mit einem Schichtplan mit fest zugeteilten Tagen und somit Zuständigkeiten funktionierte. Die engmaschige Kommunikation aller am Projekt beteiligter Moderatorinnen und Moderatoren mit den Workflowingenieurinnen und -ingenieuren stellte sich hierbei als Schlüssel für das Gelingen des Projekts und die rasche Beantwortung der Fragen und Anregungen der Freiwilligen heraus.

Zentrale Anlaufstelle hierfür ist der Diskussionsbereich des Projekts.¹⁸ Die Gestaltung von Struktur und Umfang dieses Bereichs sollte also bestenfalls bereits bei Projektbeginn überlegt und entsprechend eingerichtet werden. Zooniverse macht hierbei nur minimale Anforderungen. Die Erfahrung hat gezeigt, dass eine Unterteilung der diversen Kommunikationsbedürfnisse sehr sinnvoll ist. Um die Beiträge übersichtlich zu halten und die Anliegen voneinander zu trennen, hat sich die Aufteilung in Notes (allgemeine Kommentare zu den jeweils indexierten Dokumenten), Troubleshooting (allgemeine Fragen und Probleme) sowie Introduce yourself (hier können sich die Freiwilligen vorstellen und untereinander vernetzen) bewährt.

¹⁶ <https://www.zooniverse.org/projects/arolsen-archives/every-name-counts/talk/3279/1457620?comment=2366054&page=1>

¹⁷ Zahlen entnommen aus: <https://www.zooniverse.org/projects/arolsen-archives/every-name-counts/stats>

¹⁸ <https://www.zooniverse.org/projects/arolsen-archives/every-name-counts/talk>

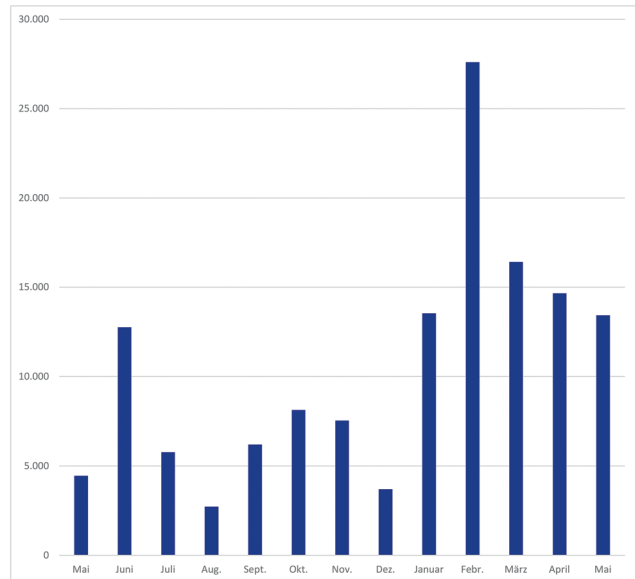


Abb. 4: Anzahl der Kommentare von Mai 2020 bis März 2021.

Seit dem Start des Projekts wurde und wird der Betreuungsbedarf der Freiwilligen immer größer. Neben denjenigen, die nur hin und wieder ein Dokument indexierten und sich im Forum kaum beteiligten, bildete sich rasch ein Kern sogenannter „Superuser“ heraus. Diese besonders engagierten Freiwilligen mit den unterschiedlichsten Hintergründen und Muttersprachen, eigneten sich innerhalb weniger Monate ein so großes Fachwissen an, dass sie die Moderatorinnen und Moderatoren der Arolsen Archives seither im Forum unterstützen und auf Fehler in den Workflows hinweisen.

„For some cards I check the address on Google maps to make sure I got the correct city and typed correctly. Seeing the building from which someone got taken makes me so aware this was just few years ago and we have to raise our voice against similar sentiments.”¹⁹

Überraschend war für uns die Menge an Daten, die wir durch #EveryNameCounts nicht nur durch die Workflows, sondern auch durch das Diskussionsforum erhielten. Die Freiwilligen begannen, ausgehend von den individuellen Dokumenten aus den Workflows, zusätzliche Informationen zu den jeweiligen NS-Verfolgten zusammenzutragen und im Notes-Thread zu verlinken. Wir erhielten somit nicht nur die abgefragten Metadaten, sondern auch unzählige Weblinks, Fotos und Geo-IDs, die für unsere Da-

¹⁹ <https://www.zooniverse.org/projects/arolsen-archives/every-name-counts/talk/3279/1459125?comment=2368308>

tennachbereitung eine zusätzliche Herausforderung darstellen. Bis wir eine Lösung für die Nutzbarmachung dieser Daten gefunden haben, speichern wir die mühevollen und für unsere Forschungs- und Erinnerungsarbeit äußerst wertvollen Rechercheergebnisse ab. Nicht nur jeder Name, sondern auch jeder Kommentar/Beitrag zählt.

Datenaufbereitung. Ein Ausblick

Der Umfang der erfassten Daten ist einerseits durch die dreifache Eingabe und die Anzahl der indexierten Dokumente beträchtlich, und wird andererseits durch die beschriebenen zusätzlichen Daten um eine nicht strukturierte und auch nur schwerlich strukturierbare Datenmenge erhöht. Die Aufbereitung der Daten nach Abschluss der Workflows und erfolgreichem Download stellt einen mehrstufigen Prozess dar, der für die Institution ebenfalls völlig neu entwickelt werden musste. Hierzu ist der Aufbau gänzlich neuer Prozessabläufe notwendig, um die Daten in das Archivinformationssystem zu importieren und letztlich im Online-Archiv bereitzustellen. Der Aufwand an personellen und technischen Ressourcen darf hierbei nicht unterschätzt werden. Auch die Unterstützung durch externe Dienstleister, denen die Besonderheiten bzgl. der Anforderung an Archivdaten zumeist nicht geläufig sind, kann den erforderlichen Zeitbedarf nur bedingt minimieren. Die Arolsen Archives befinden sich inmitten dieses Weiterentwicklungsprozesses der archivischen Datenverarbeitung.



Ramona Bräu
Manager Digital Business Development
Arolsen Archives
International Center on Nazi Persecution
Große Allee 5–9
34454 Bad Arolsen
ramona.braeu@arolsen-archives.org
www.arolsen-archives.org

Ramona Bräu ist Manager Digital Business Development bei den Arolsen Archives – International Center on Nazi Persecution. Sie studierte Geschichte und Politikwissenschaft an den Universitäten Jena und Wrocław. Nach ihrer Tätigkeit als wissenschaftliche Volontärin bei der Gedenkstätten Stiftung Buchenwald und Mittelbau-Dora war sie als Mitarbeiterin im Rahmen des Bundesforschungsprojekts zur Geschichte des Reichsministeriums der Finanzen tätig.



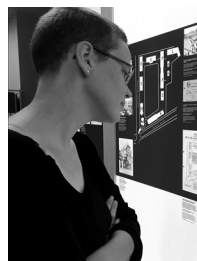
Dr. Kerstin Hofmann
Forschung und Bildung
Arolsen Archives
International Center on Nazi Persecution
Große Allee 5–9
34454 Bad Arolsen
kerstin.hofmann@arolsen-archives.org
www.arolsen-archives.org

Dr. Kerstin Hofmann ist wissenschaftliche Mitarbeiterin in der Abteilung Forschung und Bildung der Arolsen Archives – International Center on Nazi Persecution. Sie studierte Geschichte und Politikwissenschaft an der Universität Mannheim und wurde dort 2017 über die strafrechtliche Aufarbeitung von NS-Verbrechen in der Bundesrepublik promoviert. Das Projekt #EveryNameCounts begleitet sie seit Sommer 2020 als Moderatorin und Ansprechpartnerin für wissenschaftliche Fragen.



Sonja Nilson
Archiv
Arolsen Archives
International Center on Nazi Persecution
Große Allee 5–9
34454 Bad Arolsen
sonja.nilson@arolsen-archives.org
www.arolsen-archives.org

Sonja Nilson ist wissenschaftliche Mitarbeiterin der Abteilung Archiv der Arolsen Archives – International Center on Nazi Persecution. Sie studierte Geschichte und Politikwissenschaft an der Universität Düsseldorf. Sie befasst sich bei den Arolsen Archives mit Datenbanken und Datenmanagement und begleitet das Projekt #EveryNameCounts als Datenkuratorin.



Christa Zwilling-Seidenstücker
Forschung und Bildung
Arolsen Archives
International Center on Nazi Persecution
Große Allee 5–9
34454 Bad Arolsen
christa.seidenstuecker@arolsen-archives.org
www.arolsen-archives.org

Christa Zwilling-Seidenstücker ist wissenschaftliche Mitarbeiterin in der Abteilung Forschung und Bildung der Arolsen Archives – International Center on Nazi Persecution. Sie studierte Politikwissenschaften in Heidelberg und Jerusalem. Bei den Arolsen Archives ist sie Mitarbeiterin im Projektteam von #EveryNameCounts und bereitet aktuell die Schulkampagne zum Projekt vor, die im Herbst 2021 starten wird.