

From Archive to Dataset: AI-Assisted Mass Digitisation of Historical Photographic Collections, Using the Sven Simon Archive as a Case Study

Paper for the 19th Image and Research Conference in Girona, 19-20 November 2026

Responsible author: Frank Golomb; prepared with assistance from AI United Archives GmbH, Cologne, Germany
ORCID: none

Abstract

Historical photographic archives contain considerably more information than the comparatively small number of images traditionally selected for digitisation and publication. United Archives is therefore developing a scalable process that combines high-throughput scanning with AI-assisted metadata creation in order to digitise complete analogue photographic archives as comprehensively as possible. The objective is not merely to preserve individual images, but to transform entire collection contexts into structured, searchable datasets. The first large-scale implementation is the Sven Simon Archive, comprising approximately 1.6 million photographs relating to German and international political, cultural, economic and sporting history since the early 1960s.

This paper describes the collaboration with a specialised scanning company in Cologne, the linking of physical objects to digital files, the capture of inherited captions and the use of multimodal AI to produce structured metadata drafts. Experience to date shows that AI substantially accelerates descriptions, keywords, translations and the structuring of heterogeneous source information. It does not, however, replace archival contextual knowledge or the verification of people, places, dates, events and rights.

Comprehensive digitisation is particularly valuable because it reveals series, developments over time and previously overlooked relationships, while enabling new forms of data use in journalism, research, education, historical analysis and future AI applications. Prerequisites include traceable metadata provenance, risk-based quality assurance and a clear separation between source evidence, machine inference and human-verified statements.

Keywords: photographic archives; comprehensive digitisation; artificial intelligence; metadata; Sven Simon; mass scanning; data value; historical analysis

1. Introduction

Large historical photographic archives are simultaneously rich in information and poor in visibility. Their value lies not only in well-known individual images, but also in the density and temporal depth of entire collections. Nevertheless, traditional digitisation projects have generally selected presumed highlights: prominent figures, well-known events or technically outstanding photographs. This is

understandable in view of limited budgets, but it comes at a cost. The unselected portion of an archive remains effectively invisible and is available neither to researchers nor to new digital methods of analysis.

This logic of selection dates from a time when scanning, cataloguing, storing and distributing each individual image entailed substantial manual costs. It transfers analogue scarcity into the digital world. In press photography archives, however, ostensibly incidental material can be especially revealing: sequences recorded before and after a well-known event, changes in public spaces, recurring actors, behind-the-scenes workflows or the gradual evolution of visual conventions. Selection based on present-day knowledge predetermines which parts of the past will remain available for later analysis.

United Archives preserves more than forty photographers' estates and press photography archives containing many millions of originals. Its total holdings comprise approximately 25 million analogue photographs, of which around 800,000 images have been digitised to date. Since 2008, the company has focused on acquiring and safeguarding historical photographic archives, often together with documented provenance and exploitation rights. Practical experience shows that physical preservation alone is insufficient. A collection becomes a usable digital resource only when images, inherited information, context and rights are reliably connected.

Against this background, United Archives is developing an infrastructure for the most comprehensive possible digitisation of self-contained collections. "Comprehensive" does not mean that all approximately 25 million objects across the company's holdings will be processed immediately. Rather, it denotes the strategic decision to digitise selected archives as coherent bodies of evidence instead of extracting only presumed highlights from them. The first major project to follow this approach is the Sven Simon Archive, comprising approximately 1.6 million photographs.

This paper examines three questions: How can a large photographic collection be scaled to an industrial level without destroying its archival order? Which metadata tasks can AI perform, and where do human decisions remain indispensable? What additional data value is created when an archive becomes digitally available not merely in extracts, but within its temporal and thematic context?

2. The Sven Simon Archive as a Historical Dataset

2.1 Scope and significance

The Sven Simon Archive comprises approximately 1.6 million photographs documenting political, cultural, economic and sporting events in Germany and internationally since the early 1960s. Its scale makes it an important visual record of German post-war history. Its significance, however, derives not only from individual well-known photographs, but equally from several decades of continuous journalistic observation and the breadth of subjects covered.

Such a collection is more than the sum of its images. It contains chronological sequences, variants, preceding and subsequent situations, alternative perspectives and photographs that were never used editorially. These elements can demonstrate how an event was covered photographically, which images an editorial team selected and which motifs remained in the archive. For historical research,

therefore, not only what is depicted is relevant, but also the production and selection history of the photography itself.

If such a collection is searched only for prominent names or well-known events, digitisation reproduces the earlier editorial selection. By contrast, the most comprehensive possible capture opens a large "window onto the past". It makes it possible to study historical developments not only through isolated icons, but through dense visual evidence. This applies, for example, to changes in political staging, sporting culture, media work, fashion, mobility, architecture and public space.

2.2 What "comprehensive digitisation" means in practice

Comprehensiveness is not a self-explanatory concept. In a large archive, duplicates, contact sheets, technically unsuccessful photographs, different prints from the same negative and existing digital files may coexist. The objective must therefore not be confused with indiscriminately reproducing every physical manifestation.

For this project, comprehensiveness initially means considering the inherited collection context as a whole. Series and folders are not excluded merely because their present market value is unclear. Decisions about which physical objects constitute independent units of information to be digitised are based on their relationship to the photographic work, their captions, their format and their function within the archive. Duplicate detection and series formation are intended to identify redundancy without deleting it in an uncontrolled manner.

This approach differs from the unrealistic ambition of comprehensively cataloguing all 25 million United Archives photographs in the short term. Instead, the infrastructure is being developed using a defined major collection and will subsequently be transferred to other suitable archives. Comprehensiveness is therefore a strategy applied collection by collection, not a claim of unlimited capacity.

2.3 Series order as information in its own right

In press photography, a substantial proportion of the information lies between the individual images. Photographers work in sequences: they approach a situation, change position and focal length, respond to gestures and produce variants for different editorial purposes. If only the subsequently published image from such a sequence is digitised, a visual icon may be preserved, but the process by which it was created disappears.

The order of a series may show which action preceded a well-known moment, how long a situation lasted or whether a seemingly spontaneous photograph was staged. Technically weaker images can also provide contextual information, for example about people present, spatial relationships or the sequence of an appointment. This information is relevant to historical research, visual studies and media history, even if the individual image has no immediate licensing value.

The digital representation must therefore be capable of storing relationships between photographs. These include shared physical filing, consecutive numbers, chronological proximity, visual similarity

and editorial association. AI can group visually similar images and suggest possible boundaries between series. The original order must not, however, be replaced by a new, purely algorithmic order. Machine-generated clusters are enrichments; the inherited structure remains a source in its own right.

This is precisely where the additional value of comprehensive digitisation becomes apparent. With a small selection, the context of a series exists only in the memory of those familiar with the physical archive. With coherent capture, it becomes machine-readable, searchable and verifiable. The result is not simply a larger image collection, but historical data of a different quality.

3. From the Physical Archive to a Digital Process Chain

3.1 Preparation and high-throughput scanning

Digitisation is being carried out jointly with a specialised scanning company in Cologne. This collaboration separates two tasks that often take place at a single workstation in smaller projects: technically standardised image capture and archival or semantic cataloguing. The scanning company can focus on throughput, colour reproduction, resolution, cropping, orientation and file integrity. United Archives is responsible for selection rules, collection context, object assignment, the metadata model and final acceptance.

Before scanning, the material is divided into logically and physically manageable units. Folders, envelopes, sequences, captions and existing numbering systems are recorded. This preparation is not a secondary task. If context is destroyed at this stage, even powerful AI cannot reliably reconstruct it later. The Federal Agencies Digital Guidelines Initiative likewise treats technical image capture as a controlled process with measurable quality objectives (Federal Agencies Digital Guidelines Initiative, 2023).

Each transfer unit requires a unique identifier. These identifiers link the physical location, scan file, source text and subsequent metadata record. A filename alone is insufficient. Series must remain reconstructable even if individual files are rescanned, moved or transferred to other output systems. Technical quality controls verify completeness, assignment, orientation and processability before content enrichment begins.

3.2 Capturing source information

Historical photographs carry information not only within the image. Reverse sides, envelopes, index cards, agency stamps and editorial notes may contain names, dates, places, publication details, photographer attributions and rights information. This information is captured photographically or as text and assigned to the relevant object or series.

Optical character recognition can convert typewritten captions and printed agency text into machine-readable form. Handwriting, damaged labels, multilingual information and historical abbreviations nevertheless produce errors. OCR is therefore treated as derived information, not as an infallible

source. The original view is retained so that employees and future researchers can verify the transcription.

The compilation of input data is decisive for subsequent AI processing. An image alone provides a visual description. The same image combined with text from the reverse, the folder label, chronological proximity and collection affiliation enables substantially more precise cataloguing. The context must, however, be supplied to the model in separate fields so that a weak OCR reading does not silently become an apparent fact about the image.

3.3 Structured datasets instead of isolated files

The result of digitisation is not merely a directory of image files. A structured record is created for each relevant unit of information, with links to the original, scan, series, source caption, photographer, rights and processing status. This structure enables different subsequent uses without overwriting the underlying evidence.

A distribution partner, for example, requires a concise international description and controlled keywords. Academic research additionally requires source transcriptions, dating information and sequence context. A rights-management process requires photographer attribution, contractual status and restrictions. The underlying data remain the same; their output profiles differ.

3.4 Layers of the data model

The different types of information are stored in separate layers. The first layer describes the physical object and its position within the holdings. This includes the archive, sub-collection, container, format, identifier and, where appropriate, condition. The second layer contains technical information about the digital surrogate, such as filename, resolution, colour profile, creation date and quality control. The third layer preserves inherited information such as reverse-side captions, stamps, numbers and editorial notes.

Only the fourth layer contains derived metadata: normalised names, standardised dates, translations, controlled keywords and AI suggestions. A fifth layer records human review, correction and approval. This separation prevents a later improvement to the description from overwriting the historical source. At the same time, it permits an image to be reprocessed with a better model without losing the existing processing history.

For practical purposes, it is particularly important that not every image needs to receive all layers in full immediately. A minimum record may consist of an object identifier, scan, collection affiliation, source text and a verified basic description. Further enrichment is added according to the available sources, risk and intended use. This creates a tiered cataloguing model: being comprehensively digitised does not necessarily mean that every object is described to the maximum extent on the first day.

This principle resolves an apparent contradiction between volume and diligence. The technical and structural capture can encompass an entire collection, while the depth of content can vary.

Particularly relevant or sensitive images receive intensive review; other images are made discoverable at a reliable minimum level. In this way, the archive as a whole remains accessible without claiming an unaffordable, identical treatment of every metadata field.

4. AI-Assisted Metadata Creation

4.1 AI tasks

Multimodal AI is used to generate structured metadata drafts from the image, source text and context. Suitable tasks include an objective initial description, the identification of visible objects and actions, keyword suggestions, the structuring of inconsistent legacy texts, linguistic normalisation and translation. Brief agency notes can be turned into readable draft descriptions while the original transcription remains unchanged.

The machine thereby reduces the time spent confronting an empty record. Employees do not have to compose every sentence from scratch, but instead review, correct or reject suggestions. For clearly captioned and visually unambiguous images, the work shifts from writing to verification. This increases potential throughput and improves formal consistency across large volumes.

Output is generated field by field, rather than as free-form prose that merely appears finished. A typical draft contains a title, a description of what is visible, names of people taken from sources, objects, activities, suggested dates and locations, keywords and indications of uncertainty. Observation, transcription, normalisation and inference remain separate.

4.2 Limitations of image recognition

The greatest risks do not concern general object keywords, but historically consequential claims. A model may confuse a face with a prominent person, infer the wrong place from architecture or identify an ordinary press conference as a famous event. Such errors are dangerous because they sound plausible and appear to improve discoverability.

Personal identities therefore require documentary evidence or a human-confirmed identification in a controlled context. Facial similarity may provide candidates for review, but must not automatically publish a name. The same applies to precise dates, locations and events. A model can provide clues; the final statement must be based on reverse-side text, series context, contemporary documentation or reliable external sources.

AI also tends to overdescribe. A smiling person is said to be "celebrating", a crowd becomes "supporters", and a damaged facade is presented as the result of a specific conflict. Such formulations turn what is visible into claims about intention or cause. Archival descriptions, by contrast, should remain with observable actions unless a source supports further statements.

The study by Abele et al. (2025) on the use of a vision-language model in archival cataloguing confirms the importance of human review. OCR errors and hallucinations limit quality; trust is created not by linguistic fluency, but by transparent and verifiable processes.

4.3 Provenance of machine-generated enrichments

AI-generated metadata must not merge indistinguishably with historical source information. Later users must be able to determine whether a date appeared on the object, was imported from an agency database, was inferred by a model or was confirmed by a person.

Fields should therefore record their origin, time of creation, system or model version, applicable rules, review status and, where appropriate, an uncertainty category. The provenance approach of the Europeana Data Model offers an important model: enrichments are attributed to the process or actor that created them (Europeana, 2021). The IPTC Photo Metadata Standard also shows a move towards documenting AI systems in image metadata, although its new fields mainly concern AI-generated images (International Press Telecommunications Council, 2025).

Provenance has direct operational value. If the model or prompt changes, older records can be selected specifically for sampling or reprocessing. If a description is challenged, its basis can be reconstructed. Correction thereby becomes a normal component of controlled data maintenance.

4.4 Recurring errors in historical collections

In addition to incorrect identifications of people and events, several less conspicuous but systematic errors occur. Models tend to describe historical images in contemporary standard language. They may recognise a telephone, factory or demonstration, for example, but do not readily capture why the particular form is characteristic of a specific period. Historically distinctive institutions or official titles are sometimes replaced with modern equivalents. This makes a text linguistically smoother but historically less accurate.

Another problem is the resolution of ambiguity. Old captions may contain abbreviations, uncertain spellings of names or contradictory dates. A generative model often chooses a plausible variant instead of leaving the conflict visible. From an archival perspective, however, the uncertainty itself is important information. Conflicting sources must be documented alongside one another until a reliable resolution is possible.

Biases in training data also have an impact. Prominent Western individuals and places are more likely to be suggested than lesser-known actors or regions. Men in public roles may be recognised more readily than women whose names are sometimes absent from historical captions. Non-Western clothing or social contexts are occasionally generalised. Automated enrichment can therefore amplify imbalances in both the historical archive and the model.

Finally, models do not always distinguish reliably between press photography, advertising photographs, film stills and documentary scenes of everyday life. These image types may appear visually similar but have different evidential value and legal conditions. The image type must therefore be determined from collection context, markings and rights status, and must not be inferred solely from the visible scene.

These experiences lead to a simple rule: the system should not make as many claims as possible, but distinguish as clearly as possible between knowledge, source, observation and supposition. A shorter, verifiable record is more valuable than a detailed but speculative description.

5. Quality Assurance and a Changing Division of Labour

5.1 Risk-based review

Uniform, comprehensive review of every field would largely eliminate the productivity gain; unreviewed publication, by contrast, would create metadata debt. Records are therefore routed according to risk. General visible terms such as "car", "microphone" or "snow" can be checked quickly. Named individuals, exact places and dates, events, photographer attributions, rights and sensitive characteristics require more rigorous review.

Machine confidence is only one signal. Generative systems can be wrong while expressing great linguistic certainty. Routing therefore also takes into account the available sources, agreement between independent pieces of information, collection type, lists of known individuals, chronological consistency, rights status and intended channel of use. Historically significant or commercially relevant images receive increased attention.

5.2 New roles

The introduction of AI changes the division of labour more than it changes headcount. Production requires clear responsibility for the material flow, the interface with the scanning partner, file imports and exceptional cases. Metadata quality requires a separate area of responsibility with the authority to define rules, conduct sampling, stop problematic series and feed corrections back into the process.

Experienced cataloguers do not become superfluous. Their knowledge is needed to understand source structures, identify implausible outputs, develop controlled vocabularies and decide when uncertainty must remain visible. Reviewing fluent but occasionally incorrect texts demands sustained scepticism. Task rotation, sampling and escalation rules are therefore components of quality assurance.

5.3 Measurement

The relevant productivity metric is not the number of descriptions generated, but the number of approved, usable records per working hour. Preparation, failed assignments, correction and reworking must be included. Useful measures include correction time per field, the proportion requiring expert review, rejection rate, serious errors, search success and the proportion of digitised images that become available for defined uses.

For the Sven Simon project, the initial results presented here are primarily qualitative operational experience. Reliable quantitative statements about total throughput, error rates and cost per approved record require longer production runs documented in a comparable manner. A projection based on model speed has deliberately been avoided. AI can generate thousands of drafts in a short time, yet overall throughput remains dependent on preparation, scanning, assignment and review.

Experience to date already demonstrates four stable effects. First, AI significantly accelerates initial description and structuring. Second, the bottleneck shifts to review and the handling of exceptions. Third, reliable object and series identifiers become more important. Fourth, productivity is improved not by generating the greatest possible number of fields, but by limiting the metadata profile to genuine search and usage requirements.

5.4 Learning from corrections

Human corrections are not merely a cost; they are training and control data for the process. When captured in a structured manner, they reveal which fields, image types or periods are particularly prone to error. A corrected person, rejected location suggestion or misread caption can be assigned to an error class. Prompts, rules, vocabularies and routing thresholds can then be adjusted accordingly.

A distinction must be made between isolated errors and systematic patterns. A single typographical error does not justify a process change. If, however, a model repeatedly names the same incorrect institution, regularly confuses film characters with real people or modernises historical official titles, the underlying cause must be addressed. Quality assurance thus develops from a final inspection into a learning feedback loop.

Samples must not be limited to randomly selected, mostly uncomplicated images. They must be stratified by period, section of the collection, source quality, type of subject and risk class. A record with correct colours and objects must not statistically offset an incorrect personal identification. Serious errors require greater weighting than stylistic corrections.

6. The Additional Value of Comprehensive Collection Data

6.1 Research and journalism

Comprehensive collection data expands research beyond already familiar names and events. Users can find series, secondary subjects and alternative perspectives. An editorial team obtains not only the frequently published key image, but the visual context of an event. This creates new opportunities for historical photo essays, documentaries and thematic publications.

Comprehensive capture also prevents commercial value from being defined exclusively by previous demand. Subjects that appeared incidental when photographed may become highly relevant decades later. Digitisation does not merely preserve known significance; it keeps open the possibility of future reassessment.

6.2 Research, education and historical analysis

For research and education, the greatest added value arises from density, continuity and comparability. A comprehensively captured collection can reveal changes over years or decades: political gestures, media staging, sporting events, clothing styles, technical devices, urban scenes or the presence of social groups.

Image analysis must not be confused with objective historiography. Even a large press photography archive is the product of journalistic commissions, geographical priorities and editorial decisions. Comprehensive digitisation does not eliminate these biases, but makes them more open to investigation. Patterns of selection, frequencies and gaps themselves become subjects of research.

Thematic datasets can be created for different user groups without physically processing the collection again. A dataset may concern political change, football culture, media history or public architecture, for example. The crucial point is that such derivatives refer back to a traceable overall collection.

6.3 Future AI applications

Extensive, structured image collections with clarified rights may in future be used for further machine applications: similarity search, series formation, duplicate detection, entity linking, temporal analysis or domain-specific models. Their value nevertheless depends on data quality, provenance, rights and documented restrictions on use.

A historical image collection is not automatically a suitable training dataset. Incorrect identities, biased selection and unclear rights would be multiplied through machine reuse. Today's investment in traceable metadata is therefore also a prerequisite for responsible future use of AI.

6.4 Analytical potential and methodological caution

The possibility of analysing large volumes of images by machine must not be confused with a claim that the results are neutral. Frequencies in a press photography archive initially show what was photographed, acquired, captioned and preserved. They represent social reality only indirectly. A growing number of photographs on a topic may indicate actual significance, increased media attention, changing commissioning policies or better preservation.

Comprehensive digitisation nevertheless improves the conditions for investigating such biases. In a curated selection, selection decisions have already been built into the dataset and often can no longer be reconstructed. A self-contained digital collection makes it easier to compare commission density, series lengths, missing periods and editorial selection.

Analyses should therefore be based on documented research questions and disclose their data basis. Which parts of the archive were captured? Which objects were treated as duplicates? Which metadata derive from historical captions and which from AI? What uncertainties remain? Such information is not an inconvenient limitation, but a prerequisite for large visual datasets to be taken seriously in scholarship.

This approach also offers opportunities for education and public engagement. Learners can examine not only a canonical image, but also variants and context. They can understand how photographic selection shapes historical memory. The archive thereby becomes not merely a stock of illustrations, but a resource for data and media literacy.

7. Organisational and Economic Consequences

The comprehensive digitisation of a large archive is not a self-contained individual project, but the construction of a permanent data infrastructure. Models, vocabularies, customer requirements and historical knowledge change. Records must be versionable and capable of renewed enrichment without losing the original source information.

The economic value does not arise solely from the number of files scanned. The decisive question is whether those files become discoverable records that can be placed in their legal context and used across different channels. Model costs are only one part of the calculation. Material preparation, scanning, storage, OCR, integration, review and error correction may constitute larger cost blocks.

At the same time, the benefits are distributed over a longer period and across several markets. The same record can support journalism, direct licensing, partner distribution, research, education, thematic products and future analytical methods. This changes the economic logic: digitisation is not merely a cost centre for individual image sales, but an investment in reusable data assets.

The approach is transferable to other archives if four conditions are met: the physical collection must be sufficiently ordered or reconstructable; rights and provenance must be documented; scanning and data processes require shared identifiers; and quality assurance must form part of production planning from the outset. AI alone does not create scalable archive digitisation. It delivers value within a controlled process chain.

Not every collection is equally suitable. Severely disrupted orders, unclear rights or missing object relationships can increase costs substantially. Economic benefits also differ. A commercial press photography archive requires different output profiles from a municipal archive or museum. What is transferable, therefore, is not every detailed decision, but the fundamental architecture of technical capture, a preserved source layer, machine enrichment, documented provenance and risk-based review.

A further limitation lies in technological dependence. Models and providers change; interfaces, prices and data-protection conditions may also change. The archive must therefore not store its knowledge exclusively in an external system. Source information, controlled vocabularies, approval rules and processing history must be retained in a proprietary, exportable data structure. AI is a replaceable component of the infrastructure, not its memory.

The NIST recommendations on artificial intelligence risk management can be applied here: responsibilities are governed, the collection and potential harms are identified, performance is measured at the relevant level, and recognised risks are managed continuously (National Institute of Standards and Technology, 2023). Europeana similarly emphasises the importance of expertise, collaboration, appropriate training data and integration into existing infrastructure (Europeana Network Association, 2022).

8. Conclusion

The Sven Simon Archive demonstrates why the future of photographic heritage cannot lie solely in digitising selected outstanding images. Approximately 1.6 million photographs represent not merely a stock of potentially marketable individual subjects, but a coherent visual dataset relating to German and international history since the early 1960s. Its value arises from individual images, series, repetitions, developments over time and gaps in journalistic observation.

High-throughput scanning creates the technical files. AI accelerates description, structuring, keywording and translation. Yet only unambiguous object relationships, preserved source information, traceable provenance and human review turn those files into reliable historical records. The appropriate comparison is not with an autonomous archivist, but with a very fast assistant possessing a broad vocabulary, uneven historical knowledge and a tendency to conceal uncertainty.

Experience to date confirms the feasibility of a scalable approach while also revealing its limitations. The bottleneck does not disappear; it moves from writing to preparation, assignment, review and the management of exceptions. Productivity must therefore be measured by approved records, not by the number of machine-generated texts.

Comprehensive digitisation is not an absolute claim that all 25 million United Archives photographs will be captured immediately. It is a methodological decision to treat suitable archives within their inherited context and not to restrict their future through the premature selection of presumed highlights. The Sven Simon Archive serves as the first large-scale implementation and as a basis for transferring the method to other collections.

The long-term benefits extend beyond collection preservation and image licensing. Coherent digital collections enable new forms of journalism, research, education and historical analysis. They can produce thematic datasets and form the basis for future AI methods. The prerequisite is that speed and responsibility be designed together. Only then can an archive become a dataset without historical evidence becoming machine-generated invention.

References

Abele, L., Anders, G., Aydın, T., Buder, J., Fischer, H., Kimmel, D., & Huff, M. (2025). *ArchiveGPT: A human-centered evaluation of using a vision language model for image cataloguing* [Preprint]. arXiv. <https://doi.org/10.48550/arXiv.2507.07551>

Europeana. (2021). *Provenance of metadata enrichments and translations in EDM: External provenance*. https://pro.europeana.eu/files/Europeana_Professional/Share_your_data/Technical_requirements/EDM_profiles/EDM_provenance_profile_external_202111.pdf

Europeana Network Association. (2022). *AI in relation to GLAMs task force report*. https://pro.europeana.eu/files/Europeana_Professional/Europeana_Network/

Europeana_Network_Task_Forces/Final_reports/
AI%20in%20relation%20to%20GLAMs%20Task%20Force%20Report.pdf

Federal Agencies Digital Guidelines Initiative. (2023). *Technical guidelines for digitizing cultural heritage materials: Third edition*. Library of Congress.
https://www.digitizationguidelines.gov/guidelines/FADGI%20Technical%20Guidelines%20for%20Digitizing%20Cultural%20Heritage%20Materials_3rd%20Edition_05092023.pdf

International Press Telecommunications Council. (2025, November 27). *IPTC Photo Metadata Standard adds new properties for AI-generated content*. <https://iptc.org/news/iptc-photo-metadata-standard-2025-1-adds-ai-properties/>

National Institute of Standards and Technology. (2023). *Artificial intelligence risk management framework (AI RMF 1.0) (NIST AI 100-1)*. U.S. Department of Commerce.
<https://doi.org/10.6028/NIST.AI.100-1>

UNESCO. (2021). *Recommendation on the ethics of artificial intelligence*.
<https://www.unesco.org/en/legal-affairs/recommendation-ethics-artificial-intelligence>