

EuroCulture: A Multi-Strategy Approach to Large-Scale Cultural Heritage Dataset Creation

Julian Götzinger, Huajian Liu, Waldemar Berchtold, Martin Steinebach; Fraunhofer SIT | ATHENE; Darmstadt, Germany

Abstract

Cultural artifact classification supports preservation, scholarship, and anti-trafficking enforcement, but computational approaches face severe data scarcity: limited labeled collections, inconsistent metadata, and complex annotation requirements. We present a methodology for large-scale dataset creation from aggregated heritage platforms that combines cultural marker extraction from unstructured text descriptions, cross-lingual querying across multiple European languages, visual clustering using self-supervised features, and collection-aware data splitting. Applied to Europeana, this produced datasets ranging from 33,859 to 310,252 objects across 6 to 62 classes of varying granularity. Visual clustering removed 34% of noisy content while enabling fine-grained subcategories. A k -nearest-neighbor probe of the resulting feature space revealed that 79.5% of nearest neighbors share the same institutional source. This nearly matches the 83.0% same-civilization rate, indicating that institutional signatures are almost as dominant as cultural content. We quantify this bias per civilization and propose collection-aware splitting as a necessary countermeasure for realistic evaluation.

Introduction

Cultural artifacts provide important insights into humanity's past. They reflect artistic practices, technological developments, and social structures across civilizations. Effective classification of such artifacts (determining place of origin, time period, and cultural context) is valuable for scholars, conservators, law enforcement combating illicit trafficking, and cultural institutions [1, 2]. A representative use case arises at border crossings, where customs officers intercept suspicious objects and must determine geographic origin and age to assess whether an offense has occurred. Such assessments typically require expert knowledge that is rarely available on site. Computational approaches have emerged to support analysis and classification at scale, assisting non-experts in initial provenance assessments [3, 4] and enabling large-scale analyses of cultural collections [5, 6].

Data Scarcity Challenges

Progress in computational heritage is hindered by fundamental data scarcity [7]. Cultural artifacts are limited in quantity, with many lost to time or dispersed across institutions [8, 9]. Access is frequently restricted by institutional policies, preservation concerns, and copyright [10, 11, 12]. Furthermore, accurate annotation requires specialized domain knowledge spanning art history, archaeology, and cultural studies [13, 14], making large-scale labeling prohibitively expensive [15].

Existing datasets tend toward narrow geographic or temporal scopes: Examples include the AlgHeritage collection (20,000 Algerian monuments) [16], ImageOP (1,613 Brazilian religious

buildings) [17], and Karnataka inscriptions (1,748 objects) [18]. Meanwhile, aggregation platforms like Europeana contain over 50 million digitized objects, yet the vast majority remain computationally inaccessible due to inconsistent metadata standards and multilingual descriptions that limit discoverability [19]. When filtering for records with complete structured metadata in culturally relevant fields, pre-16th century dating, and image availability, fewer than 0.1% of Europeana's records met our criteria for cultural artifact classification.

Contributions

This paper presents a methodology to overcome these limitations through the following contributions:

1. A multi-strategy approach combining text mining for cultural markers, cross-lingual querying, and visual feature clustering to create large-scale heritage datasets from aggregated platforms with inconsistent metadata.
2. Multiple dataset variations from the Europeana platform, ranging from 33,859 to 310,252 objects across 6 to 62 classes.
3. Quantitative evidence of institutional collection bias via k -nearest-neighbor analysis of the visual feature space, and a collection-aware splitting methodology to enable realistic evaluation.

Related Work

Researchers have developed several approaches to mitigate data scarcity in cultural heritage, though most remain limited in scale. Custom dataset creation efforts include the AlgHeritage dataset for Algerian monuments [16], ImageOP for Brazilian churches [17], CULTURE3D for 3D landmarks [20], and Art500K for artworks [21, 22]. These are valuable but focus on narrow geographic or temporal domains.

Data augmentation and transfer learning have shown promise for small heritage datasets [18, 23], but do not tackle the underlying problem of dataset creation from heterogeneous archives.

Our group has pursued both classification and monitoring of cultural artifacts. Bugert et al. [4] and Liu et al. [24] applied transfer learning and self-attention-based feature extraction to classify artifacts by provenance, but found that museum collection taxonomy as a proxy for geographic origin limited labeling accuracy. Liu et al. [25] developed automated monitoring of stolen artifacts on online marketplaces using dedicated artifact features. Both lines of work depend on well-labeled training data, motivating the dataset creation methodology presented here.

Collaborative aggregation platforms like Europeana provide large-scale access but inherit metadata inconsistencies from con-

tributing institutions and present language barriers that limit discoverability [19]. Recent advances in vision-language models (VLMs) offer alternative approaches to heritage indexing through text-image retrieval [26], but using text-conditioned representations for both retrieval and quality assessment risks conflating textual and visual evidence. Despite these approaches, few comprehensive methodologies exist for creating large-scale, multi-cultural datasets from heterogeneous archives while accounting for the systematic biases inherent in institutional collections.

Methods

Our methodology tackles data scarcity through four complementary strategies, each targeting a specific limitation of heritage data. An overview is shown in Figure 1.

Cultural Marker Extraction

Rather than relying on structured metadata fields, which are frequently empty or inconsistent across institutions, we perform full-text search across all textual fields of Europeana records: titles, descriptions, concepts, and subject tags.

To illustrate, a limestone canopic jar in the Wellcome Collection has empty fields for date, place, type, and culture, yet its free-text description reads “Limestone human-headed canopic jar, ... from a mummified Egyptian body. This jar represents Imsety, one of the four sons of Horus.” The cultural information is present but trapped in unstructured text.

Using regex-based pattern matching, cultural markers were extracted from these text fields: nouns and adjectives indicating cultural context (e.g., “pharaonic”, “mummified”, “Horus” → “ancient Egyptian”). The extracted markers were manually reviewed to eliminate false positives and organized into civilization classes. For each class, queries were constructed by combining markers with logical OR, matching any record whose text contains at least one relevant marker.

The Europeana API was queried with minimal restrictions, focusing on images while filtering out objects labeled with modern centuries. This text-mining approach yielded 184,676 objects across 8 civilization classes from an initial 33,859 achievable through structured metadata alone, a fivefold increase. The precision of the cultural marker extraction was not formally measured against ground truth; the visual clustering step described in the following subsection serves as a complementary quality filter, removing approximately 34% of text-mined content that was visually unsuitable for classification tasks (e.g., documents, photographs of exhibition settings, and drawings).

Cross-Lingual Expansion

Many smaller European institutions describe their collections exclusively in their national language, making them invisible to English-only queries. To capture these collections, the cultural marker queries were translated into nine additional European languages: German, Spanish, French, Greek, Italian, Dutch, Polish, Romanian, and Swedish.

Each translated query was used to retrieve additional objects from Europeana. The nine languages were selected based on their observed prevalence in Europeana’s non-English metadata at the time of data collection. After deduplication, this strategy discovered 125,576 unique additional objects (approximately 40% cross-language overlap in the raw results) and expanded the

dataset from 184,676 to 310,252 objects across 16 civilization classes.

Visual Feature Extraction and Clustering

Text mining inevitably introduces noise: modern photographs, scanned documents, exhibition posters, and other non-artifact content that happens to contain cultural keywords. We apply visual clustering to identify and remove such content while simultaneously creating fine-grained subcategories.

Feature Extraction

Visual features are extracted using DINOv2-Large with registers [27] (dinov2_vitl14_reg), a Vision Transformer (ViT-L/14) with 14×14 pixel patches. Input images are resized to 256 pixels, center-cropped to 224×224 , and normalized with ImageNet statistics. For each image, we concatenate two complementary representations: the CLS token (1024 dimensions) for global semantic information and the mean-pooled patch tokens (1024 dimensions) for fine-grained local details. The resulting 2048-dimensional feature vector is then used for clustering.

We chose DINOv2 over text-conditioned alternatives such as CLIP [26] because, as a self-supervised model, DINOv2 clusters images based on purely visual similarity without textual conditioning. Using a model whose representations are influenced by text descriptions would risk circular reasoning, since those descriptions were used to retrieve the images in the first place. DINOv2 thus provides an independent visual signal for quality assessment.

Clustering and Manual Inspection

K-means clustering is applied within each civilization class, producing 2–10 clusters per class depending on class size (109 clusters total across 16 civilizations). Each cluster is manually inspected by reviewing centroid images and random samples. Clusters are then either retained and labeled with a descriptive subcategory (e.g., “ancient Egyptian vessels”, “medieval coins”) or removed if they contain non-artifact content.

Of 109 clusters, 62 were retained and 47 removed. The removed clusters comprised photographs of modern scenes (19), scanned documents (9), line drawings and illustrations (9), and other non-artifact content (10). This cleaning reduced the dataset from 310,252 to 204,909 objects across 14 civilizations (−34%), while the retained clusters define 62 fine-grained visual-cultural categories. Figure 2 shows representative examples of retained and removed clusters.

Collection-Aware Splitting

A critical challenge in heritage datasets is institutional collection bias. Objects from the same museum share systematic similarities in photography setup, lighting, background, digitization equipment, and resolution, all due to institutional provenance rather than cultural characteristics. Figure 3 illustrates this: three rows of Imperial Roman artifacts from different museums exhibit notably different visual signatures despite sharing the same civilization label.

Standard random splitting allows these signatures to leak across training and test sets, inflating evaluation metrics.

We implement group-based splitting that treats each museum collection as an indivisible unit. All objects from a given in-

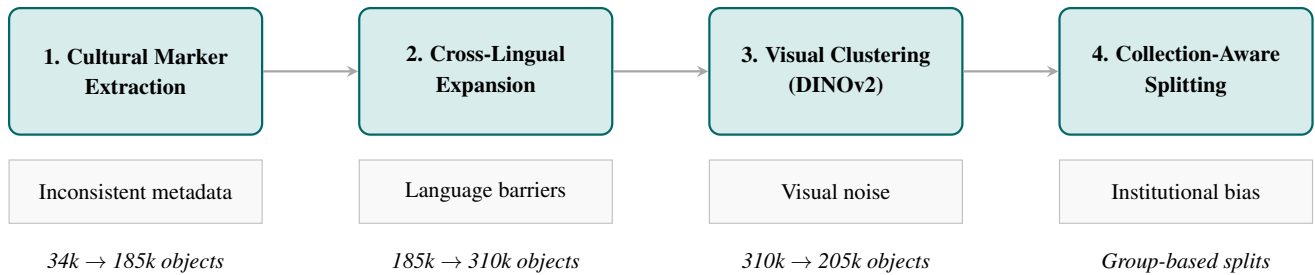


Figure 1: Pipeline overview. Each strategy addresses a specific challenge in heritage data, progressively refining the dataset.



Figure 2: Cluster quality control examples from manual inspection. Of 109 clusters, 62 were retained (a) and 47 removed (b).

stitution appear in exactly one split. While targeting 80-10-10 train/validation/test proportions, we prioritize institutional separation over exact percentages and ensure collection diversity in validation and test sets. This forces models to generalize to entirely new institutional contexts. In realistic deployment, for example, a seized artifact at a border crossing will come from an institution not represented in the training data.

Results

Dataset Overview

The multi-strategy approach produced several dataset variations with progressively increasing scale and refinement, summarized in Table 1.

Table 1: EuroCulture dataset variants.

Dataset	Objects	Classes	Strategy
S	33,859	6	Strict metadata
L	184,676	8	+ Text mining
XL	310,252	16	+ Cross-lingual
XL-clean	204,909	14	+ Visual cleaning
XL-cluster	204,909	62	+ Fine-grained categories

The expansion from S to L demonstrates the value of full-text mining over structured metadata alone. The cross-lingual ex-

pansion from L to XL captured 125,576 additional unique objects, with French and Dutch contributing the most (approximately 38,000 each), which reflects the scale of national digitization programs in those countries.

Visual clustering then removed 34% of the XL dataset as noise while enabling fine-grained subcategorization. Collection-aware split variants (XL-clean-groupsplits and XL-cluster-groupsplits) were produced from the cleaned datasets.

Comparison with Existing Datasets

Table 2 contextualizes the EuroCulture datasets relative to existing cultural heritage datasets. Previous datasets typically focus on specific regions or artifact types, while the XL dataset spans multiple civilizations and time periods. The text-mining approach avoids dependence on structured metadata, which limits most existing efforts, and the group-based splitting methodology addresses a source of bias not previously addressed in cultural heritage dataset creation.

Institutional Bias Analysis

To quantify the extent of institutional bias in the feature space, we conducted a k -nearest-neighbor (k -NN) probe on the XL-clean dataset (204,909 objects, 2048-dimensional DINOv2 features). For 50,000 randomly sampled query objects, we iden-

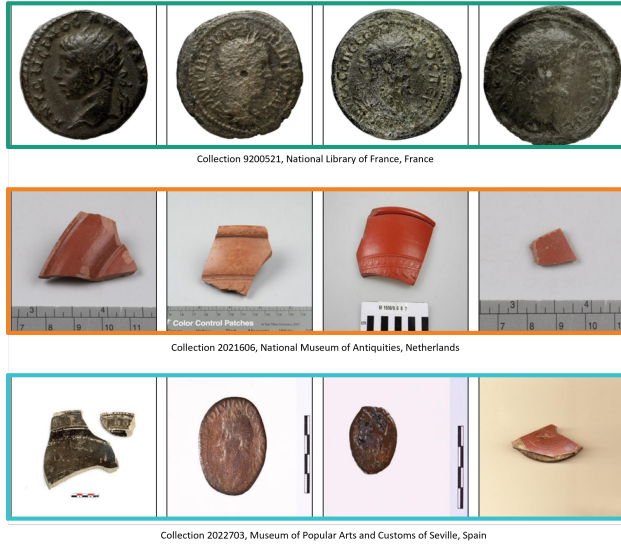


Figure 3: Institutional visual signatures within Imperial Roman. Each row: a different museum collection with distinct photography styles and backgrounds.

Table 2: Comparison with existing heritage datasets.

Dataset	Objects	Scope
AlgHeritage [16]	20,000	Algerian monuments
Karnataka [18]	1,748	Indian inscriptions
ImageOP [17]	1,613	Brazilian churches
EuroCulture S	33,859	6 civilizations
EuroCulture XL	310,252	16 civilizations
EuroCulture XL-cluster	204,909	62 visual-cultural classes

tified each object’s nearest neighbor in the full feature matrix and recorded whether it shared the same collection (institution) and/or the same civilization (class label). Results are shown in Table 3.

Table 3: k -NN probe results ($k=1$) on the XL-clean feature space.

Metric	Rate
Same collection (institution)	79.5%
Same civilization (class)	83.0%
Same collection and same civilization	73.8%
Same collection, different civilization	5.7%

The key finding is the near-equivalence of the same-collection rate (79.5%) and the same-civilization rate (83.0%), a gap of only 3.5 percentage points. This means institutional signatures explain almost as much of the nearest-neighbor structure as cultural content does. The 5.7% of cases where the nearest neighbor shares the same collection but belongs to a different civilization represents a lower bound on institutional contamination: where collection and civilization coincide, the bias is present but hidden.

The UMAP visualization in Figure 4 makes this pattern visible: the same spatial structure emerges whether points are colored by civilization or by institution, confirming that the feature space encodes institutional provenance nearly as strongly as cultural identity.

The bias varies considerably across civilizations, as shown in Table 4. Imperial Roman objects exhibit a 94.6% same-collection

rate (47 contributing collections, dominated by large numismatic archives), while Ancient Egyptian objects show only 36.0% (91 collections with more diverse institutional coverage). The range spans nearly 60 percentage points, so institutional dominance is far from uniform and depends heavily on the collection landscape of each civilization. Additionally, dataset concentration is high: on average, the top five collections per civilization account for 77.1% of all objects.

Table 4: Per-civilization k -NN bias: fraction of nearest neighbors from the same institution ($k=1$).

Civilization	Same-Coll	Colls.	Objects
Imperial Roman	94.6%	47	25,668
Stone Age	86.5%	98	20,120
Iron Age	85.0%	79	28,530
Byzantium	84.8%	92	17,308
Norse Culture	80.8%	30	1,551
Medieval	78.4%	144	38,757
Celtic Culture	77.0%	98	8,623
Bronze Age	76.8%	75	8,043
Mesopotamian	75.2%	57	4,807
Ancient Hellenic	73.8%	108	33,415
Provincial Roman	69.0%	36	6,187
Etruscan	66.7%	31	1,508
Religious Orders	65.9%	70	1,169
Ancient Egyptian	36.0%	91	9,223
Weighted average: 79.5% (291 collections, 204,909 objects)			

These results suggest that collection-aware splitting is necessary for realistic evaluation: a classifier evaluated under standard random splitting may learn institutional shortcuts instead of cultural features, producing inflated performance estimates that do not generalize to unseen institutions.

Discussion

Implications of Collection Bias

The k -NN analysis demonstrates that institutional bias is a measurable and dominant property of heritage feature spaces, not merely a theoretical concern. Classification systems trained and evaluated with standard random splits risk learning to distinguish photography setups and archival styles rather than cultural features. For practical applications such as automated provenance assessment at border crossings, evaluation results from random splitting may substantially overestimate real-world performance.

The finding that institutional signatures are within 3.5 percentage points of cultural content in determining nearest-neighbor structure (79.5% vs. 83.0%) suggests that group-based splitting should become standard practice in computational heritage. This approach reduces apparent model performance, but it provides realistic estimates of generalizability to unseen institutional contexts.

Limitations and Future Work

This work has several limitations. The methodology is built exclusively on Europeana and therefore inherits its European-centric coverage and institutional membership biases. The nine additional query languages expand coverage but exclude many European and non-European languages present in heritage collections.

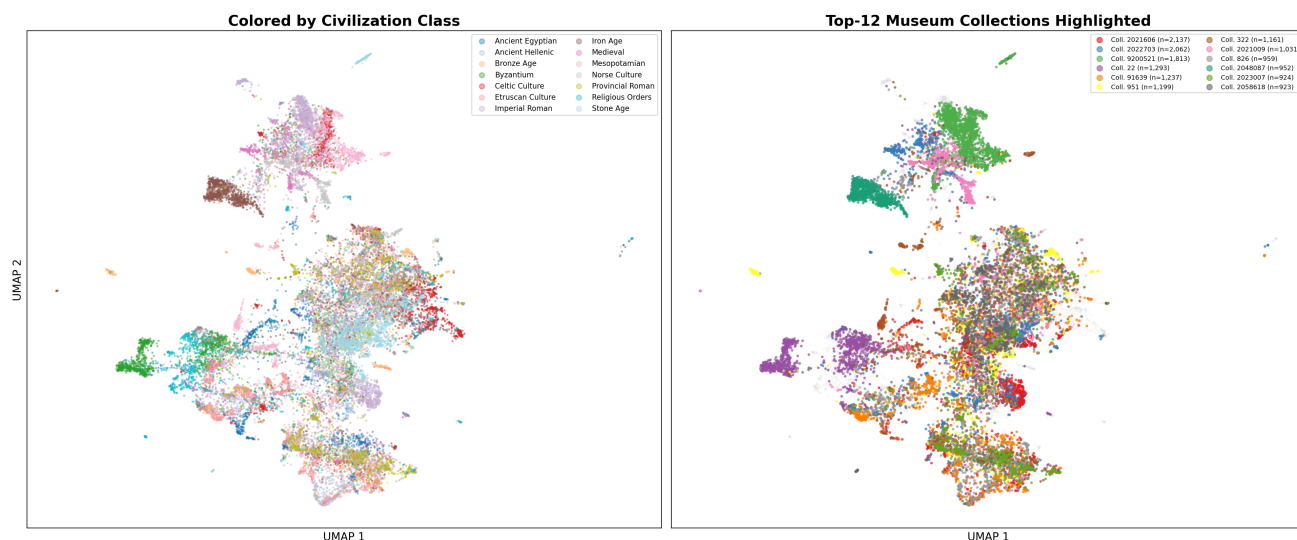


Figure 4: UMAP projection of DINOv2 features for the XL-clean dataset. **Left:** colored by civilization class. **Right:** colored by museum collection (top-12 highlighted). The visual similarity of both clustering patterns confirms that institutional signatures are nearly as dominant as cultural content in the feature space.

Manual cluster inspection is effective but does not scale: 109 clusters across 16 civilizations is feasible, but applying this approach to a broader set of civilizations or finer granularity would benefit from automated cluster quality metrics [14]. Similarly, the text-mining step uses regex-based pattern matching, which was sufficient for the current scale and offers full interpretability of matching decisions. However, this approach could be complemented by vision-language models or large language models for more robust cultural marker identification and multi-label classification, particularly for ambiguous or context-dependent descriptions [28].

The reliance on DINOv2 for visual clustering may introduce model-specific biases, as its representations are learned from general image datasets rather than heritage-specific data. Data augmentation strategies to reduce collection-specific visual signatures (e.g., background normalization or style transfer) represent an unexplored but potentially valuable direction.

Downstream classification benchmarks comparing standard and collection-aware evaluation are planned as future work.

Conclusion

This paper presented a scalable methodology for creating cultural heritage datasets from aggregated digital platforms with inconsistent metadata. By combining text mining for cultural markers, cross-lingual querying, and visual feature clustering, we transformed Europeana’s computationally inaccessible archive into structured datasets of up to 310,252 objects across 16 civilizations, subsequently cleaned and subcategorized into 204,909 objects across 62 fine-grained visual-cultural categories. A k -nearest-neighbor analysis revealed that institutional collection signatures dominate the visual feature space nearly as strongly as cultural content (79.5% same-collection versus 83.0% same-civilization nearest-neighbor rates), demonstrating that collection-aware evaluation is essential for realistic performance assessment in computational heritage applications.

Acknowledgments

This work was supported by the Federal Government Commissioner for Culture and the Media (BKM) from funds of the Federal Government’s national AI strategy, and by the National Research Center for Applied Cybersecurity ATHENE.

References

- [1] Mostafa Ibrahim. AI for Art & Heritage Conservation | Ultralytics. <https://www.ultralytics.com/blog/ai-in-art-and-cultural-heritage-conservation>, 2024.
- [2] Eva Savina Malinverni, Dante Abate, Antonia Agapiou, Francesco Di Stefano, Andrea Felicetti, Marina Paolanti, Roberto Pierdicca, and Primo Zingaretti. SIGNIFICANCE deep learning based platform to fight illicit trafficking of Cultural Heritage goods. *Scientific Reports*, 14(1):15081, July 2024.
- [3] D. Abate, A. Agapiou, K. Toumbas, A. Lampropoulos, K. Petrides, R. Pierdicca, M. Paolanti, F. Di Stefano, A. Felicetti, E. S. Malinverni, and P. Zingaretti. ARTIFICIAL INTELLIGENCE TO FIGHT ILLICIT TRAFFICKING OF CULTURAL PROPERTY. *The International Archives of the Photogrammetry, Remote Sensing and Spatial Information Sciences*, XLVIII-M-2-2023:3–10, June 2023.
- [4] Simon Bugert, Huajian Liu, Waldemar Berchtold, and Martin Steinebach. Cultural assets identification using transfer learning. *Electronic Imaging*, 34(8):273–1–273–4, January 2022.
- [5] Angela Dressen. Cultural Heritage Data for Research – An Introduction. *Open Library of Humanities*, 11(1), February 2025.
- [6] Sander Münster, Ferdinand Maiwald, Isabella Di Lenardo, Juha Henriksson, Antoine Isaac, Manuela Milica Graf, Clemens Beck, and Johan Oomen. Artificial Intelligence for Digital Heritage Innovation: Setting up a R&D Agenda for Europe. *Heritage*, 7(2):794–816, February 2024.
- [7] Simon H. Bickler. Machine Learning Arrives in Archaeology. *Advances in Archaeological Practice*, 9(2):186–191, May 2021.
- [8] Boris Jardine, Emma Kowal, and Jenny Bangham. How collections end: Objects, meaning and loss in laboratories and museums. *BJHS Themes*, 4:1–27, 2019.

- [9] Noah Charney. Lost Art: When works disappear into private collections. <https://www.theartnewspaper.com/2018/11/08/lost-art-when-works-disappear-into-private-collections>, November 2018.
- [10] Adam Buick. Copyright and AI training data—transparency to the rescue? *Journal of Intellectual Property Law and Practice*, 20(3):182–192, March 2025.
- [11] Ariadna Matas. AI ‘opt-outs’: Should cultural heritage institutions (dis)allow the mining of cultural heritage data? | Europeana PRO. <https://pro.europeana.eu/post/ai-opt-outs-should-cultural-heritage-institutions-dis-allow-the-mining-of-cultural-heritage-data>, August 2024.
- [12] OCLC Research. Making Archival and Special Collections More Accessible. OCLC Research Report, OCLC Research, Dublin, Ohio, 2015.
- [13] Luis Rei, Dunja Mladenec, Mareike Dorozynski, Franz Rottensteiner, Thomas Schleider, Raphaël Troncy, Jorge Sebastián Lozano, and Mar Gaitán Salvatella. Multimodal metadata assignment for cultural heritage artifacts. *Multimedia Systems*, 29(2):847–869, April 2023.
- [14] Anna Foka and Gabriele Griffin. AI, Cultural Heritage, and Bias: Some Key Queries That Arise from the Use of GenAI. *Heritage*, 7(11):6125–6136, October 2024.
- [15] CVAT. How Much Does it Cost to Outsource Annotation to a Data Labeling Service? | CVAT Blog. <https://www.cvat.ai/resources/blog/outourcing-data-annotation-cost>, July 2024.
- [16] Toufik Djellout and Hassina Aliane. Building and evaluation of an Algerian Cultural Heritage dataset using convolutional neural networks. In *2022 4th International Conference on Pattern Analysis and Intelligent Systems (PAIS)*, pages 1–7, Oum El Bouaghi, Algeria, October 2022. IEEE.
- [17] André Luiz Carvalho Ottoni and Lara Toledo Cordeiro Ottoni. ImageOP: The Image Dataset with Religious Buildings in the World Heritage Town of Ouro Preto for Deep Learning Classification. *Heritage*, 7(11):6499–6525, November 2024.
- [18] Bipin Nair B J, Veiy P Mwatihanye, and Akhil S. Levelwise Degradation Classification of Ancient Stone Inscriptions Images using MobileNetV2: A Deep Learning Approach. In *2024 4th International Conference on Sustainable Expert Systems (ICSES)*, pages 764–770, Kaski, Nepal, October 2024. IEEE.
- [19] Clemens Neudecker. Cultural heritage as data: Digital curation and artificial intelligence in libraries. In *Conference on Digital Curation Technologies*, 2022.
- [20] Xinyi Zheng, Steve Zhang, Weizhe Lin, Aaron Zhang, Walterio W. Mayol-Cuevas, Yunze Liu, and Junxiao Shen. CULTURE3D: A Large-Scale and Diverse Dataset of Cultural Landmarks and Terrains for Gaussian-Based Scene Rendering, 2025.
- [21] Hui Mao, Ming Cheung, and James She. Deepart: Learning joint representations of visual arts. In *Proceedings of the 25th ACM international conference on Multimedia*, pages 1183–1191. ACM, 2017.
- [22] Hui Mao, James She, and Ming Cheung. Visual arts search on mobile devices. *ACM Transactions on Multimedia Computing, Communications, and Applications (TOMM)*, 15(2s), 2019.
- [23] Mahdi Bahrami and Amir Albadvi. Deep Learning for Identifying Iran’s Cultural Heritage Buildings in Need of Conservation Using Image Classification and Grad-CAM. *Journal on Computing and Cultural Heritage*, 17(1):1–20, February 2024.
- [24] Huajian Liu, Xiaoying Yang, Raphael Antonius Frick, and Martin Steinebach. Identification of cultural artifacts using deep learning. *Electronic Imaging*, 37(8):271–1–271–7, 2025.
- [25] Huajian Liu, York Yannikos, Julian Heeger, Simon Bugert, Walde-mar Berchtold, and Martin Steinebach. Automated monitoring of stolen cultural artifacts on online marketplaces. *Electronic Imaging*, 37(3):314–1–314–6, 2025.
- [26] Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, Gretchen Krueger, and Ilya Sutskever. Learning transferable visual models from natural language supervision. In *International Conference on Machine Learning*, pages 8748–8763. PMLR, 2021.
- [27] Maxime Oquab, Timothée Darcet, Théo Moutakanni, Huy Vo, Marc Szafraniec, Vasil Khalidov, Pierre Fernandez, Daniel Haziza, Francisco Massa, Alaaeldin El-Nouby, Mahmoud Assran, Nicolas Bal-las, Wojciech Galuba, Russell Howes, Po-Yao Huang, Shang-Wen Li, Ishan Misra, Michael Rabbat, Vasu Sharma, Gabriel Synnaeve, Hu Xu, Hervé Jegou, Julien Mairal, Patrick Labatut, Armand Joulin, and Piotr Bojanowski. DINOv2: Learning Robust Visual Features without Supervision, February 2024.
- [28] Caroline Sporleder. Natural Language Processing for Cultural Heritage Domains. *Language and Linguistics Compass*, 4(9):750–768, September 2010.

Author Biography

Julian Götzinger is a research associate at Fraunhofer SIT / ATHENE Center in Darmstadt. He obtained his Master’s degrees in Computer Science and IT Security from TU Darmstadt in 2024. In the “Media Security & IT Forensics” department, he researches deep learning-based face verification methods, with a particular focus on applications for the protection of minors. His broader research interests include fairness and robustness in biometric systems, explainability and interpretability of deep learning models, as well as synthetic data generation.

Huajian Liu received his B.S. and M.S. degrees in electronic engineering from Dalian University of Technology, China, in 1999 and 2002, respectively, and his Ph.D. degree in computer science from Technical University of Darmstadt, Germany, in 2008. He is currently a senior research scientist at Fraunhofer Institute for Secure Information Technology (SIT). His major research interests include information security, digital watermarking, robust hashing and digital forensics.

Waldemar Berchtold has headed the Multimedia Security research group since 2022 at the Fraunhofer Institute for Secure Information Technology (SIT) and is a researcher at the National Research Center for Applied Cybersecurity (ATHENE) in Darmstadt, Germany. He received his diploma in mathematics in 2008 and his Ph.D. in 2022 at TU Darmstadt. The focus of his research is in various areas of multimedia security for authenticity and integrity proof and digital watermarking. He has led numerous projects in the field of media security with a focus on audio and video.

Martin Steinebach is the manager of the Media Security and IT Forensics division at Fraunhofer SIT. From 2003 to 2007 he managed the Media Security in IT division at Fraunhofer IPSI. He studied computer science at the Technical University of Darmstadt and finished his diploma thesis on copyright protection for digital audio in 1999. In 2003 he received his PhD at the Technical University of Darmstadt for this work on digital audio watermarking. In 2016 he became honorary professor at the TU Darmstadt. He gives lectures on Multimedia Security as well as Civil Security. He is Principal Investigator at ATHENE and represents IT Forensics and AI security. Before he was Principal Investigator at CASED with the topics Multimedia Security and IT Forensics.

JOIN US AT THE NEXT EI!

electronic IMAGING

Imaging across applications . . . Where industry and academia meet!



- **SHORT COURSES • EXHIBITS • DEMONSTRATION SESSION • PLENARY TALKS •**
- **INTERACTIVE PAPER SESSION • SPECIAL EVENTS • TECHNICAL SESSIONS •**

www.electronicimaging.org

