

EU AI-Act: Tagging GenAI Content

Julian Heeger, Waldemar Berchtold, Simon Bugert, Martin Steinebach; Fraunhofer Institute for Secure Information Technology SIT / ATHENE; Darmstadt, Hesse/Germany

Abstract

We introduce an initial framework for content traceability in AI-generated media, aligning with the objectives of the *EU AI Act*. The rapid advancements in generative AI (genAI) necessitate the development of reliable mechanisms for identifying and tracking AI-generated content to ensure transparency, trust and regulatory compliance. To address these challenges, we propose a conceptual infrastructure that facilitates media content registration for AI companies, artists and institutions. It enables provenance tracking and content authentication. Importantly, the proposed system is applicable not only to AI-generated content but also to non-AI-generated media. This dual functionality enhances trust beyond the requirements set forth in the *EU AI Act* by ensuring the identification of both authentic and synthetic content.

The framework incorporates *robust hashing* techniques, *digital signatures*, and a database to mitigate the spread of media with uncertain provenance while adhering to regulatory guidelines. A key component of this approach is the adoption of the ISO-standardized *International Standard Content Code (ISCC)* as a robust hashing method. The ISCC's decentralized architecture allows for independent implementation without legal constraints, and its adaptability ensures compatibility across various content formats. However, maintaining the flexibility to update hashing algorithms remains essential to address evolving technological advancements and adversarial manipulations.

Introduction

The need for clear identification of *generative AI (genAI)* content arises from the growing challenges associated with artificial intelligence systems capable of producing highly realistic synthetic material. These AI-generated outputs are increasingly difficult for human observers to differentiate from authentic, human-created content. Rapid advancements in AI technology, combined with widespread accessibility, have raised significant concerns regarding information integrity, trust, and security within digital ecosystems. Key risks include disinformation, large-scale manipulation, fraud, impersonation, cyber mobbing and consumer deception. Without reliable identification mechanisms, users remain vulnerable to these threats, which can significantly erode societal trust in digital communications.

To mitigate these risks, effective identification mechanisms must be implemented to ensure that AI-generated or AI-altered content is clearly distinguishable from original content, such as photos. Technical solutions such as digital watermarks, metadata embedding, cryptographic authentication, and content provenance tracking can aid in tracing the origin of AI-generated material while maintaining transparency. These methods must be reliable, interoperable, and resilient against adversarial manipulation and technological advancements.

The ability to mark AI-generated content serves several crit-

ical functions. It supports efforts to preserve credibility in digital information, combats misinformation, and reduces the risk of fraud and identity misrepresentation. As AI models continue to evolve, the necessity of implementing such solutions becomes more pressing. By embedding these mechanisms at the model or system level, developers and downstream providers can adhere to regulatory frameworks, such as the *EU AI Act*, ensuring that AI-generated content remains transparent and traceable. This approach not only protects consumers from deception but also fosters public trust in AI-driven digital environments.

Background

In this section, we will look at the framework conditions as set out in the *AI Act*¹. As a disclaimer, it should be said that none of the authors has a legal background.

AI Office

The *European Artificial Intelligence (AI) Office*², inaugurated in February 2024, serves as the central authority for AI regulation and development within the *European Union (EU)*. Its primary mission is to ensure the coherent implementation and enforcement of the *AI Act*, particularly concerning general-purpose AI models, while fostering an ecosystem of excellence and innovation in trustworthy AI [2].

The AI Office comprises five specialized units:

- 1. Regulation and Compliance Unit:** Coordinates the EU's regulatory approach to AI, ensuring uniform application and enforcement of the *AI Act* across member states. This unit also contributes to investigations, addresses potential infringements, and administers sanctions when necessary.
- 2. AI Safety Unit:** Focuses on identifying systemic risks associated with general-purpose AI models, developing mitigation strategies, and establishing evaluation and testing methodologies to ensure AI safety.
- 3. Excellence in AI and Robotics Unit:** Supports and funds research and development initiatives to cultivate an ecosystem of excellence in AI and robotics. It coordinates programs like *GenAI4EU*³, promoting the growth and integration of AI models into innovative applications.
- 4. AI for Societal Good Unit:** Designs and implements projects leveraging AI for public benefit, such as enhancing weather modeling, improving cancer diagnostics, and developing digital twins for urban reconstruction.

¹<https://artificialintelligenceact.eu/the-act/>

²<https://digital-strategy.ec.europa.eu/en/policies/ai-office>

³https://eic.ec.europa.eu/eic-funding-opportunities/eic-accelerator/eic-accelerator-challenges-2025/genai4eu-creating-european-champions-generative-ai_en

5. **AI Innovation and Policy Coordination Unit:** Oversees the execution of the EU's AI strategy, monitors trends and investments, stimulates AI adoption through European Digital Innovation Hubs, and supports regulatory sandboxes and real-world testing environments to foster innovation

In its enforcement role, the AI Office collaborates closely with member states and the *European Artificial Intelligence Board* to ensure consistent application of AI regulations. It also engages with developers, the scientific community, and other stakeholders to draft codes of practice, conduct evaluations of general-purpose AI models, and, when necessary, impose sanctions to uphold compliance.[3].

Beyond regulation, the AI Office promotes an innovative EU ecosystem for trustworthy AI by advising on best practices, facilitating access to AI sandboxes, and supporting structures like Testing and Experimentation Facilities, *European Digital Innovation Hubs*, and *AI Factories*. These efforts aim to stimulate investment and position the EU as a global leader in responsible AI development [3].

EU-AI Act

The *AI Act* is the European Union's regulatory framework for artificial intelligence, designed to ensure safety, transparency, and compliance with fundamental rights. It follows a risk-based approach, categorizing AI systems into four levels. Unacceptable risk AI applications, such as social scoring, are banned due to their threats to safety or rights. High-risk AI, used in critical areas like healthcare and law enforcement, is subject to strict regulations, including data governance, transparency, and human oversight. Limited-risk AI systems, such as chatbots, require basic transparency measures, such as disclosing AI-generated content. Minimal-risk AI, including general-purpose systems like recommendation engines, remains largely unregulated.

The *AI Act* also establishes market surveillance, harmonized enforcement, and an AI Office to oversee compliance. It fosters trustworthy AI while promoting innovation through regulatory sandboxes.

Scope

This EU regulation applies to all entities that develop, use, import, or distribute AI systems within the EU, regardless of their location. It also covers AI systems used in the EU, even if produced elsewhere. Exemptions include AI for military, defense, national security, law enforcement by foreign authorities, and international judicial cooperation, provided individual rights are upheld. It does not apply to AI for scientific research, pre-market systems, personal non-professional use, or open-source AI unless classified as high-risk or subject to specific provisions. Existing EU laws on data protection, privacy, and confidentiality remain unaffected.

Trace of Origin

The increasing sophistication of artificial intelligence in content generation necessitates the development of robust mechanisms for tracing the origin of the content. To address these challenges, providers of AI systems must integrate technical solutions that allow for the identification and labeling of AI-generated or AI-modified content. Such labeling should be implemented in

a standardized, machine-readable format to ensure transparency and accountability across digital ecosystems.

Effective implementation requires careful consideration of the unique characteristics and constraints of different content types. For instance, while watermarking techniques may be effective for images and audio, they are not well established for text-based AI outputs. Furthermore, with advances in AI technology and the state of the art of each technology, to ensure that methods remain effective and resilient to adversarial circumvention. By doing so, AI providers comply with regulations and ethical AI practices and contribute to a more trustworthy digital landscape.

In addition to digital watermarking, the *EU AI Act* also highlights metadata embedding as a potential approach for content attribution. However, a fundamental limitation of metadata is its susceptibility to detachment from the associated content without altering the semantic integrity of the media. As a result, metadata alone is insufficient for ensuring the traceability of origin unless supplemented by additional technological measures.

Cryptographic authentication is another technique referenced in the *EU AI Act* within this context. While cryptographic methods provide robust integrity verification, their effectiveness is compromised when lossy transformations such as compression are applied. Consequently, cryptographic authentication alone is not a viable solution for origin tracing without the integration of complementary technical mechanisms that ensure resilience against content modifications.

State of the Art

In this section, we look at current developments within the topics addressed and briefly present another concept with a similar objective.

Project Origin by C2PA

Project Origin is an alliance of leading organizations from the publishing and technology sectors, dedicated to combating misinformation by establishing a verifiable chain of trust for digital media. The Coalition for Content Provenance and Authenticity (C2PA⁴) system establishes content provenance through structured metadata called *Assertions*, which document asset creation, authorship, edits, and other trust signals. These Assertions are combined into a digitally signed entity called a *Claim*, which can also incorporate W3C Verifiable Credentials for additional trust validation. A *Claim Generator*, either hardware or software, binds these elements into a verifiable C2PA *Manifest*, stored within the asset's *Manifest Store*. Trust decisions are based on the cryptographic identity of the signer (e.g. the identity associated with the cryptographic signing key), which may be a person, service, or trusted hardware. C2PA Manifests remain verifiable indefinitely, even if signing credentials later expire or are revoked.

A common scenario is a journalist capturing a photo with a C2PA-enabled camera or phone. The device generates a manifest containing Assertions such as camera details, a thumbnail, and cryptographic hashes linking the image to its provenance data. These Assertions are compiled into a digitally signed Claim, which is embedded into the JPEG file, ensuring the manifest remains valid indefinitely. A Manifest Consumer, like a C2PA Validator, can later verify the digital signature, validate Assertions,

⁴<https://c2pa.org/>

and present the information to users, helping them assess the trustworthiness of the digital content.

C2PA does not specify or recommend specific trust lists or public key infrastructures (PKI). Instead, it treats them as configurable inputs. Each application using C2PA operates within its own unique ecosystem with distinct trust requirements.

International Standard Content Code

The ISCC is an open, decentralized identifier designed to universally and uniquely identify digital content across various media types, including text, images, audio, and video. Unlike traditional identifiers that require centralized assignment, the ISCC is generated algorithmically directly from the content itself, enabling decentralized issuance and ensuring that identical content yields the same ISCC, regardless of location or context. As the name suggests, ISCC is an internationally standardized procedure in ISO/IEC 24138[1].

ISCC involves four different data sources, namely metadata, content, data, and instance.

Metadata Processing: Basic metadata associated with the content is processed to generate a Meta-Code, capturing essential descriptive information.

Content Analysis: The actual content undergoes analysis to produce a Content-Code. This involves extracting and normalizing the content, followed by applying similarity-preserving hashing algorithms to capture the content's unique characteristics.

Data Encoding: The raw data of the content is processed to create a Data-Code, representing the binary essence of the file.

Instance Identification: An Instance-Code is generated to capture the specific instance of the content, aiding in distinguishing between different copies or versions.

These components are then concatenated with a common header to form the complete ISCC, resulting in a compact, self-describing identifier that encapsulates various facets of the content.

Decentralized Issuance: ISCCs are created directly from the content without the need for a central registration authority, allowing any party with access to the content to generate its ISCC.

Similarity Preservation: The ISCC employs content-derived, locality-sensitive hashing techniques to produce identifiers that reflect the similarity between different pieces of content. This feature facilitates efficient content deduplication, clustering, and version control.

Hierarchical Structure: ISCC is a composite identifier comprising multiple self-describing components, each representing different aspects of the content, such as metadata, semantic content, and raw data. This modular design supports granular identification and management of content and its components.

The ISCC's open specification and reference implementations are publicly available, encouraging widespread adoption and integration into diverse digital content management systems.

Watermarking

Digital watermarking is a technique used to embed imperceptible or perceptible information within digital media, such as images, audio, video, text or documents, to assert ownership, ensure authenticity, or enable traceability [4], [6]. The embedded watermark can be later extracted or detected to verify the integrity or origin of the media content.

Digital watermarking offers several advantages, including copyright protection by allowing content creators to assert ownership, content authentication by verifying data integrity, and forensic tracking by embedding unique identifiers to trace unauthorized distribution. Additionally, watermarking is resistant to common transformations like compression or resizing, making it useful for broadcast monitoring, AI-generated content labeling, and sometimes used in deepfake detection.

Despite these advantages, digital watermarking has inherent limitations. It is susceptible to attacks, including collusion, filtering, and noise addition, which may compromise its robustness. There is a trade-off between imperceptibility and resilience, as strong watermarks may degrade media quality, while imperceptible ones may be easily removed. Computational complexity can be an issue, increasing file size and processing overhead. Furthermore, standardization challenges and false positives in detection can reduce effectiveness. Ethical concerns also arise regarding user privacy and unauthorized tracking.

Probably the biggest problem with digital watermarks is their symmetrical design. This means that anyone who embeds it can also read it again and vice versa. This makes it easy for anyone to manipulate a digital watermark to subsequently delete it, to gain a person's trust by concealing its origin from an AI, or to embed a digital watermark to undermine trust by pretending that the content was generated by an AI. To our knowledge, no watermarking process has yet been able to invalidate this point.

Robust Hashing

A robust hash is a type of hash function that generates a fingerprint of digital content, designed to remain stable under minor modifications while still uniquely identifying the underlying data. Unlike cryptographic hash functions (e.g., *SHA-256*), which produce completely different outputs with even the smallest input change, robust hashes are designed to tolerate distortions such as compression, format conversion, scaling, or slight noise interference.

Robust hashing typically involves feature extraction from the content rather than a direct byte-level computation. In multimedia applications such as image, video, and audio hashing, key characteristics—such as dominant frequencies, edge structures, or statistical patterns—are used to generate a hash. These extracted features are then processed through a hashing algorithm to produce a compact representation of the content.

Common techniques for robust hashing include four different approaches like perceptual hashing, which focuses on human-perceivable features [8], [12], *locality-sensitive hashing (LSH)*, which maps similar items to close hash values [7], [5], wavelet and fourier transform-based approaches, which extract key frequency domain features [13] and blockhash, which computes the mean of the pixels in each block [10].

Unlike digital watermarking, robust hashes do not require modifications to the original data. Robust hashes can be derived

from content that has already been distributed. Hash-based comparison is computationally efficient and enables large-scale indexing and searching.

In spite of the advantages of robust hashes, there are also disadvantages. Attacks on robust hashes [9] aim to change the robust hash through various post-processing steps or by machine learning to such an extent that the content is still semantically identical, but the content can no longer be identified using the hash. Furthermore, collisions occur with robust hashes for similar content [11]. In addition, the context of the paper must take into account that the robust hash procedure is known and can therefore be optimized against it. For example, an image manipulation can be optimized so that the hash is still considered known, but the image content differs semantically from the original.

Concept

In the following section, we propose a conceptual framework for addressing the challenges posed by genAI. The underlying premise is that by leveraging our understanding of specific media files, particularly their origins, we can enhance our ability to interpret the content we consume. Our approach involves facilitating the participation of artists and media entities, while also mandating that companies generating AI content submit their media files to a central database for identification. This measure is intended to mitigate the impact of media content of uncertain origin, thereby enhancing the clarity and reliability of information exchange. It is important to note that this concept is founded on regulatory requirements rather than offering a purely technical solution, a necessity which appears to align with the EU's intentions. The architecture proposed for genAI content could also be utilized by artists and media outlets to submit their released media content, thereby facilitating the identification of the source for a specific image that is currently in circulation on social media, for example.

Our architecture, illustrated in Figure 1, follows a left-to-right sequence. It involves three distinct entities:

- (a) An AI company that offers a service to generate or modify media content using genAI.
- (b) An artist who publishes their work and indicates whether genAI was used for either partial or full creation.
- (c) Institutions that publish online media content and may serve as proxies for trusted reporters and artists. User trust in a media file depends on the institution's credibility. Reporters can remain anonymous, with institutions acting as trusted intermediaries, particularly in sensitive contexts such as crisis reporting. Similarly, anonymous artists can release their work through these institutions, analogous to case (b).

For AI companies, this requires that all generated media content be sent from their servers to the database after being presented to users. This ensures that any media a user interacts with is already labeled as AI-generated in the database, allowing consumers to identify it.

For artists and institutions, media publication primarily serves to trace content origins, such as identifying misuse of an image. Institutions can publish media with additional metadata, including permanent links to relevant articles, to clearly establish its source. This can be done at any possible period, but from a misuse perspective, the sooner, the better would be sensible.

Each of the entities is issued a certificate to sign the publication of media content and to validate the correctness of the uploaded metadata. The public key infrastructure (PKI) supporting the concept could either be a new one governed by the AI Office or delegated to the member states or an existing one.

The media file is hashed using a robust hashing algorithm, as detailed in Section , to produce a unique textual representation of its specific format. This hash is then cryptographically signed using the entity's certificate, ensuring a verifiable link between the entity and the published media file. The signature and robust hash are embedded in the file's metadata for efficient retrieval upon publication. Since metadata structures vary across formats, appropriate fields must be selected accordingly. For example, JPEG files use the Exchangeable Image File Format (EXIF) to store additional information, requiring format-specific integration decisions. To ensure integrity, authenticity, and temporal validation, a Timestamp Authority (TSA) is employed in the signature process, proving the file's existence at a specific point in time. This mechanism allows institutions and artists to assert originality and enables tracking of genAI content modifications. Following the media file processing, there are two options for database integration:

- Uploading only the signed hash along with metadata.
- Uploading both the media file and metadata.

The choice depends on the entity's privacy preferences. For artists and institutions, the option to withhold media files preserves confidentiality. However, for AI companies, mandatory media file uploads should be enforced to facilitate content validation, ensuring users can verify whether their content matches the database records. A detailed discussion on this validation process is provided in Chapter .

Content detection

Content detection refers to the process of analyzing a newly received media file from an unknown source on a device, such as a smartphone or PC, to retrieve relevant metadata. The operating system, such as Android or iOS for mobile devices, or an application on a PC, can provide interfaces to generate the robust hash and manage communication with the database.

The detection process, illustrated in Figure 2, follows these steps:

1. Compute the robust hash using the ISCC algorithm.
2. Transmit the generated hash to the database.
3. Retrieve the signed response from the database containing metadata and verification details.

The database response contains media files and their associated metadata, as detailed in Chapter , and is cryptographically signed by the database operator. This enables the use of proxy servers to enhance user privacy while ensuring the authenticity and integrity of the retrieved data. By rotating proxies for each request, users can minimize the information exposed to the database operator, preventing the accumulation of identifiable user data.

Upon receiving the response the application can present the most relevant matching media file along with its metadata. This allows users to assess whether the media file displayed represents the original source and determine whether any modifications have

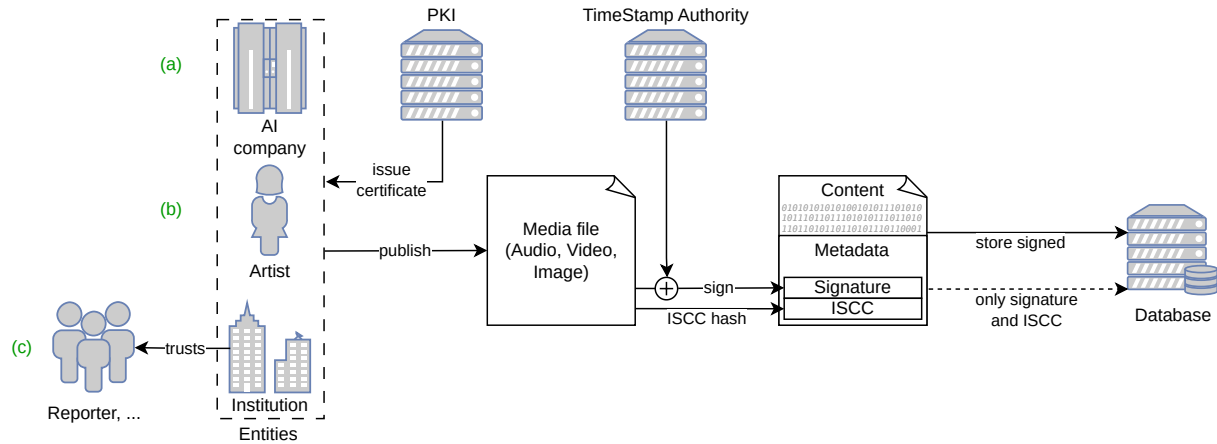


Figure 1: Architecture Overview

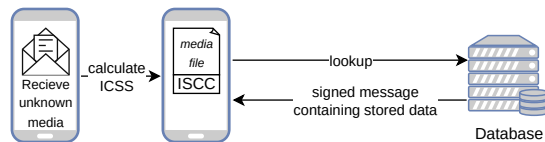


Figure 2: Lookup of unknown Media using Smartphone

been made. The accuracy of this assessment is constrained by the robustness of the hashing mechanism. Minor modifications do not affect the robust hash, ensuring that the identified media file remains linked to the original. However, significant alterations disrupt the robust hash, resulting in no match. This behavior is intentional, as substantial modifications fundamentally transform the media file, making it a distinct creation rather than a derivative of the original.

If the database response does not provide the media file, the approach depends on the intended objective:

- **GenAI-generated or modified content:** If a similar hash exists in the database and the metadata attributes the content to an AI company rather than an institution, the media file is classified as GenAI.
- **Media file origin:** The response includes metadata from the original upload, providing relevant details and references for further verification.

Discussion

Any codes of practice the AI Office will be presenting after Mai 2025, will play a decisive role in determining which technical solutions become compliant with the *AI Act* in the first place. Our presented concept is a proposal for an infrastructure to reduce the number of media content without any trace of origin. As described earlier the solution remains a regulatory rather than technical one, the implementation of which is linked to the AI office making it mandatory for companies.

The ISCC illustrates a robust hashing method suitable for adoption. It aligns with key criteria expected from a technology endorsed by the AI Office, namely being an open standard with substantial development the authors already put into it. The utilization of an open standard enables companies to implement their own versions without legal constraints and without infringing any

existing patents. However, different formats require specific algorithms and the flexibility to adapt when better alternatives emerge or existing methods prove inadequate. The system must therefore support changing the algorithm and updating the existing information in the database or at least have a migration strategy. Nevertheless, this discussion is beyond the scope of this paper and may be addressed in future publications as further guidance from the AI Office becomes available.

In Chapter , we presented an alternative approach, backed by major corporations with significant stakes, particularly those subject to the *AI Act*. Both C2PA and this approach share key principles: reliance on a public key infrastructure (PKI) for user identification and adherence to open standards to facilitate seamless integration into existing pipelines. This indicates that specific key technologies and their requirements will be important for any approach.

The database presents a distinct challenge. It could be managed either by the AI Office or an authorized entity, introducing a single point of failure. Alternatively, database management could be delegated to each individual EU member states, necessitating synchronization across all instances to ensure consistent responses to queries throughout the EU. This is fundamentally a regulatory matter that falls beyond the scope of this paper.

Conclusion and Future work

As genAI becomes more advanced, it's important to have reliable ways to trace AI-generated content, thereby ensuring transparency, trust, and compliance with regulatory frameworks such as the *EU AI Act*. This paper proposes a conceptual infrastructure that enables AI companies, artists, and institutions to register media content, facilitating origin tracing and content authentication. By leveraging robust hashing techniques, and digital signatures, this system aims to reduce the circulation of media with uncertain provenance while supporting regulatory compliance.

A key component of this approach is the integration of the ISCCs, which provides a robust hashing method suitable for all media content as an open standard. Its open nature allows for independent implementation without legal constraints, while its adaptability ensures compatibility across different content formats. However, the ability to migrate to alternative robust hash-

ing algorithms is essential to address emerging technological advances and evasion strategies employed by malicious actors.

We provide an initial concept for the traceability of content in AI-generated media, in line with the objectives of the *EU AI Act*. The final implementation will depend on regulatory requirements and guiding principles from the AI Office. These will be particularly relevant in relation to compliance standards and best practice.

Future research should focus on refining the proposed infrastructure and explore alternative robust hashing solutions.

Acknowledgments

This work has been funded by the German Federal Ministry of Education and Research (BMBF) and the Hessian Ministry of Higher Education, Research, Science, and the Arts within their joint support of the National Research Center for Applied Cybersecurity ATHENE.

References

- [1] ISO/IEC 24138. Information and documentation — international standard content code (iscc). In *Standard and International Organization for Standardization*, Geneva, CH.
- [2] European Commission. Commission Decision Establishing the European AI Office. <https://digital-strategy.ec.europa.eu/en/library/commission-decision-establishing-european-ai-office>, Jan 2024.
- [3] European Commission. The AI Office: What is it, and how does it work? <https://artificialintelligenceact.eu/the-ai-office-summary/>, Mar 2024.
- [4] I. Cox, M. Miller, J. Bloom, and M. Miller. *Digital Watermarking*. The Morgan Kaufmann Series in Multimedia Information and Systems. Morgan Kaufmann, 2001.
- [5] Mayur Datar, Nicole Immorlica, Piotr Indyk, and Vahab S. Mirrokni. Locality-sensitive hashing scheme based on p-stable distributions. In *Proceedings of the Twentieth Annual Symposium on Computational Geometry*, SCG '04, page 253–262, New York, NY, USA, 2004. Association for Computing Machinery.
- [6] J. Dittmann, A. Mukherjee, and M. Steinebach. Media-independent watermarking classification and the need for combining digital video and audio watermarking for media authentication. In *Proceedings International Conference on Information Technology: Coding and Computing (Cat. No. PR00540)*, pages 62–67, 2000.
- [7] Aristides Gionis, Piotr Indyk, and Rajeev Motwani. Similarity search in high dimensions via hashing. In *Proceedings of the 25th International Conference on Very Large Data Bases*, VLDB '99, page 518–529, San Francisco, CA, USA, 1999. Morgan Kaufmann Publishers Inc.
- [8] V. Monga and B.L. Evans. Robust perceptual image hashing using feature points. In *2004 International Conference on Image Processing, 2004. ICIP '04.*, volume 1, pages 677–680 Vol. 1, 2004.
- [9] Jonathan Prokos, Tushar M Jois, Neil Fendley, Roei Schuster, Matthew Green, Eran Tromer, and Yinzi Cao. Squint hard enough: Evaluating perceptual hashing with machine learning. *Cryptology ePrint Archive*, 2021.
- [10] Martin Steinebach. Robust hashing for efficient forensic analysis of image sets. In Pavel Gladyshev and Marcus K. Rogers, editors, *Digital Forensics and Cyber Crime*, pages 180–187, Berlin, Heidelberg, 2012. Springer Berlin Heidelberg.
- [11] Martin Steinebach. An analysis of photodna. In *Proceedings of the 18th International Conference on Availability, Reliability and Security*, pages 1–8, 2023.
- [12] Zhenjun Tang, Xianquan Zhang, and Shichao Zhang. Robust perceptual image hashing based on ring partition and nmf. *IEEE Transactions on Knowledge and Data Engineering*, 26(3):711–724, 2014.
- [13] R. Venkatesan, S.-M. Koon, M.H. Jakubowski, and P. Moulin. Robust image hashing. In *Proceedings 2000 International Conference on Image Processing (Cat. No. 00CH37101)*, volume 3, pages 664–666 vol.3, 2000.

Author Biography

Julian Heeger is a research associate in the Media Security and IT Forensics department at the Fraunhofer Institute for Secure Information Technology (SIT) and a researcher at the National Research Center for Applied Cybersecurity (ATHENE) in Darmstadt, Germany. He holds a Master's degree in IT security from the Technical University of Darmstadt.

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.

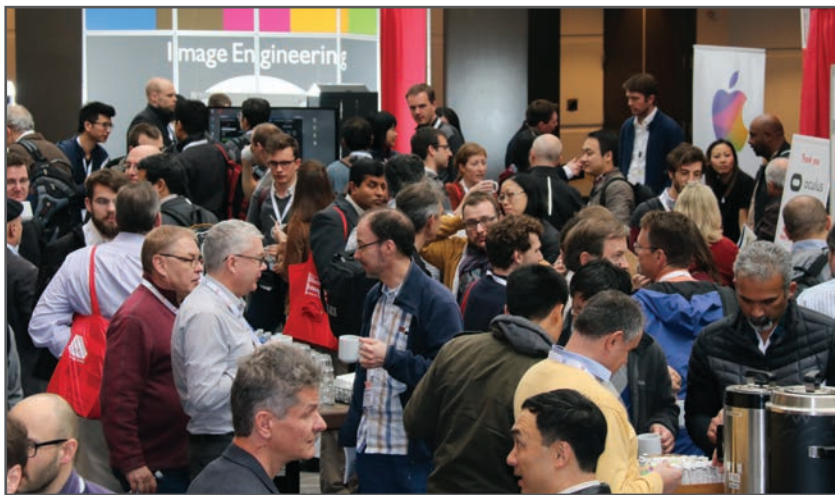
Simon Bugert received his Master's degree in computer science from the Technical University of Darmstadt, Germany in 2021. Since then, has been a research associate in the Media Security and IT Forensics department at the Fraunhofer Institute for Secure Information Technology (SIT) and at the ATHENE National Research Center for Applied Cybersecurity.

Martin Steinebach is the head of the Media Security and IT Forensics department at Fraunhofer SIT. He studied computer science and received his PhD from the Technical University of Darmstadt for this work on digital audio watermarking in 2003. From 2003 to 2007 he was head of the Media Security in IT department at Fraunhofer IPSI. In 2016 he became honorary professor at the TU Darmstadt and gives lectures on multimedia security as well as civil security. He is principle investigator at ATHENE and represents IT Forensics and AI Security. Previously, he was principle 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

