

Content ARCs: Decentralized Content Rights in the Age of Generative AI

Kar Balan¹, Andrew Gilbert¹, and John Collomosse¹

¹*DECaDE Centre for the Decentralized Digital Economy, University of Surrey, Guildford, UK*

Abstract

The rise of Generative AI (GenAI) has sparked significant debate over balancing the interests of creative rightsholders and AI developers. As GenAI models are trained on vast datasets that often include copyrighted material, questions around fair compensation and proper attribution have become increasingly urgent. To address these challenges, this paper proposes a framework called *Content ARCs* (Authenticity, Rights, Compensation). By combining open standards for provenance and dynamic licensing with data attribution, and decentralized technologies, Content ARCs create a mechanism for managing rights and compensating creators for using their work in AI training. We characterize several nascent works in the AI data licensing space within Content ARCs and identify where challenges remain to fully implement the end-to-end framework.

1 Introduction

Generative AI (GenAI) is transforming creative workflows through tools for efficiently generating and manipulating text, images, music, and video. However, the vast datasets used to train commercial GenAI models often include copyrighted material [1], causing growing concerns among creative rightsholders about how their work is used and whether they will receive proper recognition and compensation. These concerns have recently surfaced in proposals—and counter-proposals—around changes to copyright legislation in several jurisdictions. The debate has become highly contested; for example, the recent UK copyright consultation [2] received over 11,000 submissions following a public awareness campaign by creative practitioners. Arguments often centre on the contentious issue of opt-in versus opt-out mechanisms; *i.e.* whether permission should always be sought to train commercial AI models on copyrighted data (*i.e.* opt-in) or there should be a default legal right to train unless counter-indicated by a rightsholder (*i.e.* opt-out). Rightsholders often argue that an opt-in model is essential to protect creators’ control over their work and ensure fair compensation [3]. While, AI developers often contend that an opt-out approach is necessary given the sheer scale of data (often many billions of assets) over which contractual consent and payment would need to be orchestrated [4]. Political factors also weigh on the debate, such as creating a competitive AI development environment versus jurisdictions with more permissive legislation or less stringent enforcement [5]. These positions are often deeply polarized and appear difficult to reconcile, highlighting the need for balanced solutions that uphold both creative rights and technological progress.

In this paper, we contribute a framework to help resolve this deadlock by leveraging a combination of open standards and emerging data-centric technologies such as data attribution and distributed ledger technology (DLT). We present this as a meta-system design comprising of three core phases (ARC) applied to content: **(A)uthenticity**,

(R)ights, **(C)ompensation**; hereafter **Content ARCs**. Multiple technology choices exist for each phase, and we describe technical options for each.

The motivation of Content ARCs is to provide an automated mechanism for clearing rights to, and compensating rightsholders for, the use of copyrighted data in GenAI training. In this way, concerns around the practicality of opt-in are alleviated through an interoperable machine-actionable protocol leveraging open standards. While this framework may not always be necessary—AI developers may own, or collectively license [6], large datasets to support training—much of the data used for training GenAI models is drawn from the open internet or other fragmented sources where collective licensing is not feasible. In such cases, a scalable mechanism for licensing content from and returning value to individual rightsholders is essential to support large-scale AI development while ensuring fair compensation for creators. We also discuss legislative support that could benefit the adoption of Content ARCs in the creator economy.

2 Background

Technologies to express creator preferences on AI training or processing are (a) site-based (location-level) or (b) unit-based (asset-level), each approach having strengths and limitations in scope and persistence across the content supply chain.

Site-based methods allow a creator or content host to apply a blanket permission or restriction to assets on all or part of a website. For decades, the convention (RFC 9309) of parsing a file named ‘robots.txt’ on a server has served as a signal indicating which parts of a website are permitted to be scraped for search indexing. Since different crawlers can be (dis-)allowed per location, and AI training bots self-identify by name, the file can be readily applied for opt-out. Recently, the W3C has proposed a JSON-based variant (TDMRep) [7], with language aligned to the EU 2019 Copyright Directive [8], specifically Article 4, which describes an opt-out mechanism via the Text

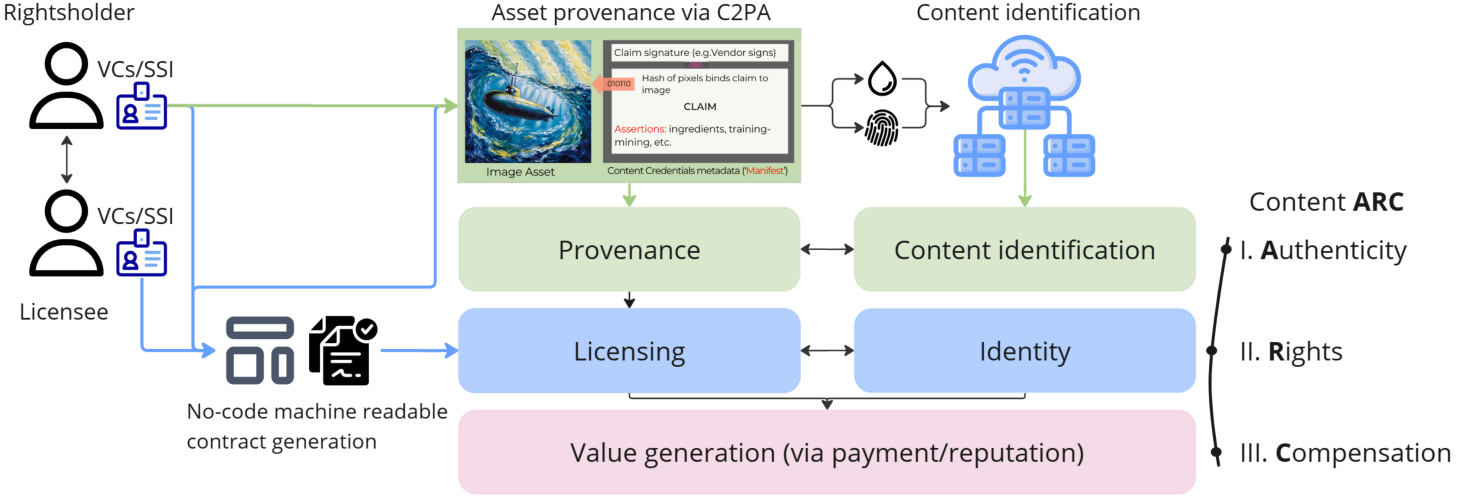


Figure 1: Content Authenticity, Rights, and Compensation (ARCs) framework, illustrating how content provenance, identification, licensing, and creator identity enable downstream value generation. Green arrows indicate systems feeding into content provenance, while blue arrows feed into licensing.

and Data Mining (TDM) exception. This use of the term ‘Data Mining’ here is interpreted as a catch-all for any AI use, reflecting the tension between rapid technology advances and the pace of legislation, which, even in 2019, was too early to consider GenAI explicitly. Ultimately, both approaches allow for efficient expression of opt-in/out in bulk; however, as the signal is not embedded in the assets, it does not persist when content is copied, shared, or aggregated downstream. This limits its use in enforcing opt-in/out within the content supply chain, and there is no mechanism for specifying licensing arrangements for AI re-use.

Unit-based methods allow creators to attach permissions and restrictions directly to individual assets by embedding machine-readable metadata within the file using standards for expressing opt-in/out. Several standards have emerged, ensuring interoperability for this task. The IPTC Photo metadata standard (2024.1) [9] includes a flag to indicate ‘Data Mining’ opt-out, which, similar to TDMRep [9], is a catch-all to indicate opt-out from any form of AI use. There is no opt-in flag, nor any implication to be drawn as to granting consent in the absence of the opt-out. Increasingly adopted is the C2PA (Coalition for Content Provenance and Authenticity) [10] metadata standard that can also be applied to assets of various modalities (e.g. image, video, audio, text) through embedding of XMP metadata within the file, or a side-car (accompanying linked file) if the format does not support metadata embedding. C2PA is primarily designed to communicate provenance information *i.e.* an audit history of how an asset is created and supports indication of opt-in/out consent within that audit trail. Furthermore, opt-in/out may be specified separately for granular AI uses (*e.g.* training or inference) and for generative, non-generative AI, or non-AI analytics (data mining). The metadata is also cryptographically signed to create a non-repudiable, tamper-evident record attached to the asset.

Content Identification. The persistence of unit-based metadata through the content supply chain is a challenge for emerging standards like C2PA that have not reached ubiquity in their adoption. Many content platforms (particularly social platforms commonly used to redistribute content) strip

such metadata from assets. Therefore, retaining a copy of metadata (including opt-in/out preferences) within a registry is common, coupled with a content identification (Content ID) scheme for looking up that metadata. Examples include invisible watermarking techniques [11, 12] that actively inject an identifier into content and content fingerprinting (perceptual hashing) techniques that passively enable identification [13–15]. DECORAIT [16] combines fingerprinting with a decentralized metadata registry held on a DLT to look up the AI opt-in/out status of images, expressed in C2PA metadata, while IPTC-PLUS [17] used fingerprinting to look up the opt-out status of images expressed in IPTC metadata.

Digital Rights Management (DRM) refers to technical mechanisms which enforce access control over digital content. Unlike opt-in/out signals, which are primarily declarative forms of usage right, DRM actively prevents unauthorized use or copying of content through encryption and other safeguards. In Content ARCs, we suggest licenses for content use and value creation and assert that such licenses should be offered legal support. We do not suggest using them as an access control mechanism, as with DRM.

3 Findings

We propose a framework called *Content ARCs* (Authenticity, Rights, Compensation) to address the growing challenges around fair compensation and proper attribution for AI-generated content. Content ARCs aim to provide a scalable, interoperable mechanism for clearing rights and compensating rightsholders for using copyrighted data in GenAI training. As illustrated in fig. 1, the framework is structured around three interconnected phases—Authenticity, Rights, and Compensation—all supported by open standards and decentralized technologies. This design allows for automated rights management and value creation at scale, even when content is distributed or aggregated across multiple platforms. We outline the key function of each phase and discuss technology choices available for their instantiation.

(A)uthenticity phase provides a verifiable way to establish

the identity and verify the authenticity of a digital asset.

(R)ights phase provides a way for a rightsholder to assert ownership of a digital asset and to create digital contracts to license that asset for use by others. Underpinned by digital identity and representations of rights.

(C)ompensation phase enables rightsholders to extract value when a licensee exercises a granted license. This includes potentially determining the level of payment via automated means *e.g.* smart contracts and data attribution technologies.

3.1 Authenticity

To license and enable value creation from media, there must first be a basis for identifying assets and establishing their authenticity [18].

The Coalition for Content Provenance and Authenticity (C2PA standard) is a cross-industry group that has developed an open standard for encoding provenance metadata directly into media files to help determine their authenticity. The C2PA standard [10] defines a data structure known as a manifest, which records key facts—called assertions—about the origin and edit history of an asset. These assertions can include details such as the tools used and any source assets (or ingredients) incorporated into the work. Ingredients themselves can carry C2PA manifests, creating a structured provenance graph. Assertions are organized into a claim, which is cryptographically signed using Public Key Infrastructure (PKI) to ensure authenticity and prevent tampering. The only mandatory assertion is a cryptographic hash of the asset which securely links the manifest to its content. By examining the provenance of an asset, an end-user may draw a conclusion regarding its authenticity. Recently, C2PA version 2 has been fast-tracked as an international standard (ISO/DIS 22144), and JPEG Trust also adopted version 1 as part of its international standard (ISO/IEC 21617-1:2025). JPEG Trust [19] is technically interoperable with C2PA and re-articulates several C2PA constructs *e.g.* a ‘Manifest’ is referred to as ‘Trust Manifest’. It introduces adjunct constructs such as the ‘Trust Profile’, a JSON-based schema for checking if assertions in an asset match some validation criteria. It may be used to help automate authenticity decisions. Several other metadata standards (IPTC, EXIF) exist for embedding general information at the time of creation, although these are not explicitly focused on authenticity, nor are they cryptographically signed. C2PA has seen strong adoption in the image modality with several camera manufacturers, digital tools (such as Adobe Photoshop and Microsoft Designer) and most Generative AI models (such as DALL-E and Adobe Firefly), and some early adoption in the video space [20].

Metadata schemes provide a way to embed a unique identifier (GUID) in an asset, for example: a Manifest ID in C2PA; or EIDR, for IDs in audiovisual content. Critics of metadata schemes, point to their fragility due to metadata stripping on content platforms (c.f. section 2), motivating the need to ground the identifier in a durable identifier derived from content itself (‘content ID’), such as a watermark or fingerprint. In contrast to the open nature of metadata standards, content ID schemes are often proprietary. This presents a challenge to ARCs in enabling a low-friction, in-

teroperable solution for robust decentralized rights management. Recently, the International Standard Content Code (ISCC) [14] became an international standard (ISO 24138), describing within a standard method to encode image fingerprints. However, the method adopted is a decades-old perceptual hash (pHash [21]) that exhibits a high false positive matching rate [15] at global scale (*i.e.* over trillions of assets). This is due to the simplistic encoding method based on frequency analysis (DCT) which cannot discriminate smaller differences in images, and its small (64 bit) key space. The cryptographic ‘birthday attack’ estimates the probability of a collision in n bits as $P(n) \approx 1 - e^{-\frac{n^2}{2^{65}}}$ *i.e.* a 50% chance at only 5 billion assets for 64 bits. Other open technologies based on DCT are PhotoDNA [22] (144 bits, mainly used for CSAM identification) and PDQ [23] (256 bits) however these still produce false matches under small manipulations. Modern, performant methods for content ID rely upon a trained neural network (AI) to derive content ID *e.g.* [13, 15, 24]. Herein lies a further paradox. Open standards are the only practical solution to adoption and interoperability. However the pace of standardization is at odds with fast moving emerging technologies like content ID; and functionality defined by a trained AI model is difficult to standardize. Similar arguments may be made for watermarking, where older methods based on frequency domain processing are being superseded by robust AI approaches [11]. This is leading toward de facto standardization practices where open sourced content ID models are being commercially adopted (*e.g.* TrustMark [11]), creating silos of interoperable solutions.

Registries play a critical role in this phase by linking a content ID (the key), to a copy of authenticity metadata *e.g.* C2PA or IPTC (the value). A centralized global registry (‘key-value store’) is impractical, although registries do exist in specific verticals (like EIDR for advertising). A federated registry structure with a unified query point could allow AI developers to identify and verify the authenticity of content before using it for model training. Distributed Ledger Technology (DLT), colloquially referred to as ‘blockchain,’ provides a way to create a decentralized key-value store without relying upon the trust of a single actor or centralized oversight. The scalability of DLT has improved with more recent proof-of-stake (PoS) networks delivering lower cost and latency.

3.2 Rights

The rights phase defines the legal and contractual terms under which content can be used, and requires a non-repudiable mechanism for a rightsholder to issue licenses to use an asset.

Some metadata standards covered within subsec. 3.1 contain a rudimentary rights declaration to indicate opt-in/out of AI asset use. IPTC has a Data Mining opt-out field. C2PA has a Training and Data Mining assertion that can indicate opt-in/out of different processes covering inference vs. training, and generative vs. non-generative AI use¹. Yet none of these standards are aimed at comprehensive rights representation, and a more general rights description language would be required for even more granular permissions (*e.g.* specific people, or sector-specific uses). As such, Content ARCs leverage open rights expression standards to communicate these terms

¹In C2PA v2 the AI opt-in/out assertion was moved from the core specification into a community extension called CAWG.

in a machine-readable format. Existing standards include the use of W3C RDF by Creative Commons [25], and the ODRL (Open Digital Rights Language) [26]. Unlike traditional static licensing models, ODRL-based licenses can be updated dynamically as licensing terms change. This implies the need for a further mechanism, external to the asset itself, to track the issuance and revocation of licenses in a non-repudiable way.

One such mechanism is the use of digitally signed license files. These could be distributed alongside the asset itself or stored externally in a registry. A digitally signed file provides cryptographic proof of authenticity and integrity, ensuring that the terms of the license are tamper-evident and traceable to the issuing rightsholder. A DLT-based licensing registry could be coupled with ODRL to allow rightsholders to update licensing terms dynamically, including revocation or conditional licensing (e.g., time-limited or use-limited licenses). Much as with content ID use cases, a DLT based registry provides a decentralized, transparent, and tamper-evident mechanism for representing license issuance.

Digital identity is a key consideration of ownership and licensing, vouching for a physical person as a digital token. Many centralized identity systems exist, often grounded in government identity document (or 'know your customer' practices). Decentralized identity schemes gaining traction are verifiable credentials and decentralized identity. Yet the systems remain technically complex and not widely adopted by end-users, including in the creative sector. Since rights systems are often coupled with DLT registries, the public wallet address of the DLT is often used to double as a user identifier.

While provenance has seen increasing standardisation through initiatives like C2PA, no equivalent effort yet exists for machine-readable rights and compensation. We highlight the need for coordinated standardisation across these phases to enable end-to-end interoperability and scalable, decentralised licensing for content.

3.3 Compensation

The Compensation phase addresses how creators are compensated for the use of their content in AI training and downstream value creation. Compensation models in Content ARCs are based on licensing terms defined in the Rights phase and enforced through automated or manual payment mechanisms. In contrast to historic Digital Rights Management (DRM) solutions that gate access, Content ARCs seek to empower creators through flexible licensing relying upon contractual terms for enforcement. Banks of standardized, vetted contracts could in the future reduce friction for creative practitioners. For example, a creator could select a pre-configured contract for AI training, specify terms (e.g., duration, permitted uses). Compensation models are coupled to value in the content supply chain, which remains under-researched in the age of generative AI [27], but may include:

Royalty or Event based models. Smart contracts automate royalty distribution whenever a licensed work is accessed, replicated, or used in AI pipelines. For instance, NFTs that represent licenses and encode licensing terms in associated SCs can trigger a micropayment to the content owner each time a dataset is downloaded or a derived asset is sold [18].

Attribution-Driven Payouts. The training datasets responsible for creating AI models can be described via provenance metadata attached to models, raising the opportunity to pay royalties to contributors for use of the model *e.g.* when synthetic assets are generated [18]. However, for Generative AI the scale of such datasets runs into the billions, presenting a challenge to fractional compensation. Data attribution can help identify the subset of training data most influential in a given synthetic asset, to target payout. Several techniques have measured attribution as visual correlation between output and training data [28, 29]. Yet correlation does not equal causation. Causative methods have been proposed, such as invisible watermarking in ProMark [30] to measure which mixture of watermarks exists in the output, as well as Shapley value-based analytics [31], however neither scales to billions of data items. Scaling causative attribution techniques to the billion-scale datasets used for training GenAI models remains an open research challenge in the development of royalty mechanisms for GenAI.

Non-Financial Incentives. Beyond direct monetary remuneration, compensation may include access to AI tools, promotional visibility (*e.g.* search ranking reward), future licensing discounts or community incentives.

4 Discussion

We discuss how the landscape of nascent decentralized content licensing implementations map onto the Content ARCs framework and summarize relevant systems in table 1.

Non-Fungible Tokens (NFTs). There has been a growing trend in the art world to represent physical art as NFTs, offering a digital tokenized version to facilitate trading and verification of ownership. These NFTs are often accompanied by different types of certificates of authenticity. This smart-contract-based trading approach also enables royalties from secondary sales, providing greater control over intellectual property rights and revenue streams, enhancing transparency and fairness for creators. Verisart (www.verisart.com) verifies artist identity through 3rd party government-issued ID checks and registers digital certificates of authenticity for both physical and digital art on the Bitcoin blockchain. Arcual (www.arcual.com) issues digitally signed certificates of authenticity and integrates smart contracts to automate licensing agreements and royalty distribution, ensuring artists receive compensation from secondary sales. Artory (www.artory.com) partners with experts to verify artwork provenance and register cryptographically signed records, additionally incorporating public auction data. Artclear (www.artclear.com) fingerprints physical artworks using microscopic imaging which records unique surface characteristics. The resulting digital fingerprints, along with associated metadata, are immutably stored on a blockchain. Comparing new scans with registered fingerprints enables the verification of the artwork's authenticity. NFT platforms rarely address the topic of copyright or rights assignment. Some platforms offer downstream compensation for resale of assets within the same platform, but downstream compensation for general re-use (including for AI training) is unaddressed.

EKILA and the ORA Framework [28] is an early technical framework and instantiation of mechanisms that aligns

Method	I. Authenticity		II. Rights		III. Compensation	
	Content ID	Verification	Representation	Identity	Attribution	Value Exchange
EKILA (ORA) [28]	C2PA soft binding (fingerprinting and/or watermarking).	Cryptographically signed provenance (C2PA).	NFTs for licenses expressed in natural language.	Ethereum wallet address.	Proportionate attribution via fingerprint for downstream compensation.	Crypto-currency micropayment via SC.
Ocean Protocol	Not implemented at the unit (asset) level.	Not implemented.	Data NFTs (ownership) + Datatokens (access rights as ERC-20 sub-licenses).	Ethereum & EVM compatible network wallet address.	Not implemented.	Datatokens (ERC-20) via SC.
Story Protocol	Not implemented. Supports watermark asset specified in metadata.	JSON metadata file and Proof of Creativity (IP provenance graph).	IP asset as NFT (ownership) + License Tokens as NFTs (licensing agreements).	Story wallet address.	Derivative works tracking and fractional royalties distribution through License Tokens.	Royalties distributed via SC in native IP token.
Vana Protocol	Not implemented.	Attestations for data quality, but authenticity is not considered.	Tokens represent fractional ownership and governance of DataDAO.	Vana wallet address.	Not implemented.	Distributed via SCs in native VANA token, but only for top 16 DataDAOs.
SongBits	Not implemented.	Not implemented.	NFTs represent shares of royalty rights.	SUI wallet address, no additional guarantees for artist identities.	Not implemented.	Distributed via SCs in native SUI token.
JPEG Trust [19] Draft v2	C2PA soft binding (fingerprinting and/or watermarking).	Cryptographically verifiable provenance information through the Trust Profile (JSON-based schema).	Open digital rights language (ORDL) and Trust Manifest checking. Rights registry.	Verifiable Credentials / DIDs (CAWG).	Not implemented.	Not implemented.
Fox Verify	Cryptographic hashing and fingerprinting.	Cryptographically signed provenance data (non-standard).	Licenses are implemented as logic within SCs.	Custom identity registry SC links cryptographic key pairs to real-world identities.	Partial implementation via ContentGraph and perceptual hash, but no automated downstream compensation.	License sales via SC in MATIC (Polygon DLT) token, no downstream royalties.

Table 1: Existing creative rights management and monetization systems mapped to the Content ARCs framework across its three core phases: Authenticity, Rights, and Compensation. Smart contract is abbreviated as SC. Yellow indicates component is present in solution. Gray indicates partially present. White indicates absent.

with all of Content ARCs but does not fully deliver across all components. It enables recognizing and rewarding data contributors to GenAI model training using a combination of the C2PA standard and NFTs stored on the Ethereum DLT to create a technical framework known as ORA (Ownership-Rights-Attribution). ORA instantiated the authenticity phase of Content ARCs by leveraging C2PA to record verifiable provenance of digital assets, which in turn can leverage C2PA soft bindings (watermarking and/or fingerprinting) to identify content. By design, C2PA communicates neither the ownership of an asset, nor (beyond the ability to express opt-in/out permissions for AI) usage rights or licensing. Rights and ownership are therefore represented as NFTs, creating an on-chain, machine-readable licensing system. Ownership is expressed by minting the C2PA-signed asset as a regular NFT. Token-based licenses are issued and managed via smart contracts operated by the rightsholder. However ORA lacks standardized rights representation, specifying only that license detail may be expressed as free-text with the NFT. For compensation, ORA hard codes an API to issue royalty payments for content reuse (in particular, for GenAI) through micropayments into a stored value system enabled via smart contracts. Given a synthetic image created by a GenAI model, EKILA applies data attribution using a correlation (fingerprint similarity method) to identify the subset of training images most responsible for that generated image and issue micropayment royalties to their owners.

Ocean Protocol (www.oceanprotocol.com) is a decentralized data exchange platform that leverages DLT to facilitate secure and transparent data sharing and monetization. Ocean Protocol publishes datasets as data NFTs, representing entire datasets rather than individual data points. While the protocol addresses certain aspects of the Content ARCs framework, it does not fully implement all phases. Data authenticity is not implemented—the blockchain records dataset

transactions, however provenance of individual data points is not verified prior to being ingested into the system. Data owners hold data NFTs representing copyright or exclusive license for the data asset. Access to datasets is granted through data-tokens—fungible ERC-20 tokens that function as sub-licenses, enabling monetization. Smart contracts enforce terms of use, but no attribution mechanisms are implemented, with data owners only receiving compensation when consumers purchase their datatokens for data access.

The Story Protocol (www.story.foundation) provides a DLT-based solution for managing the lifecycle of intellectual property assets, allowing creators to register, license, and monetize their works through smart contracts. Story provides solutions that intersect with each of the Content ARCs phases, albeit with distinctions and areas for further alignment. Each IP asset is represented as an NFT, establishing proof of ownership and origin since ingestion into the system. The Story Protocol enables derivative works to be linked back to the original, establishing a chain of provenance through the Proof of Creativity mechanism. The protocol supports uploading JSON metadata for an IP asset according to a specific structure, providing provenance details like creation timestamp, author, and media type. It also supports the registration and display of AI pipeline metadata, enabling tracking of relationships such as those between models and datasets (e.g., trained_on, finetuned_from). However, authenticity of assets is not verified based on this information prior to being ingested. For rights management, rightsholders define usage rights, royalties, and derivative work restrictions via smart contracts which distribute License Tokens and automatically enforce these terms. For the compensation phase, the Royalty Module facilitates smart contract-based royalty distribution, supporting complex royalty flows, such as compensating original creators when derivative works generate revenue.

The Vana Protocol (www.vana.org) is a DLT-based sys-

tem for AI training data ownership, governance and monetization via DataDAOs and tokenized incentives. While it aligns with aspects of the Content ARCs framework, limitations exist in data authenticity, individual control, and attribution. A Proof-of-Contribution mechanism assesses data quality and generates off-chain attestations linked on-chain. However, it does not verify that data’s provenance or authenticity. Data contributors add validated data to DataDAOs in exchange for tokens, which grant voting power. Data usage rights are determined collectively rather than individually. Contributors cannot define granular licenses, as governance decisions are based on majority rule within the DataDAO. Users earn tokens for contributing data and if their data is used in AI training. However, rewards are only distributed to the top 16 Data Liquidity Pools (DLPs) each epoch, introducing compensation uncertainty and encouraging centralization in larger DLPs.

SongBits (www.songbits.com) introduced a model for music ownership that directly involves fans in the financial success of artists. While it implements limited rights and compensation mechanisms within the Content ARCs framework, it lacks content authenticity or artist identity verification. Artists sell a share of their song’s royalty rights as digital ‘bits’ (NFTs), representing a fractional share of the song’s revenue. For rightsholder compensation, SongBits manages the collection of royalties generated from streaming platforms and distributes these earnings proportionally to shareholders.

JPEG Trust [19] is an international standard (ISO/IEC 21617-1) which, in version 1, addresses content provenance and authenticity. Its Media Tokenization group is drafting a version 2 of the standard that develops support for Identity and Rights Declaration on top of C2PA. JPEG Trust is compatible with C2PA, leveraging its soft binding (watermarking and fingerprint) capabilities for content-dependent identification and downstream attribution. Integrating the CAWG identity assertions, it uses decentralized identifiers (DIDs) for identity. Attribution is referred to in the context of attributing rights, rather than as fractional compensation scheme, using machine-readable rights expressed via the ODRL [26]. These are encoded within the image metadata and could therefore facilitate automated rule-based micro-licensing, though specific compensation mechanisms are not defined. Details of asset tokenization and registries (called rights exchanges) are not specified and deferred as implementation choices. This scoping provides flexibility but may come at the cost of interoperability in those phases.

Fox Verify (www.verify.fox) is a DLT-based content licensing and provenance system, that represents assets as hierarchical NFTs (EIP-6150) and uses smart contracts for rights management. It implements the Authenticity component of the Content ARCs framework, with limited support for Rights and Compensation. For the authenticity component, content identification is achieved through cryptographic hashing, generating unique identifiers for NFTs in conjunction with perceptual hashing via PDQ [23] to perform asset lookup via a ‘Verify’ tool to retrieve and view asset metadata. Cryptographic signatures, linked to an on-chain identity registry of publisher key pairs, provide a verifiable attestation of content origin and ownership. Rights management is implemented through machine-readable smart contract licenses, enabling structured access and usage control. Digital assets

are represented as NFTs within a ‘ContentGraph’ smart contract, which provides a structured and extensible framework for managing rights information. The hierarchical NFT structure facilitates license inheritance from parent to child nodes, facilitating rights management across complex content collections. Fox Verify enables content monetization through smart contract-based license sales, with transactions conducted in Polygon’s MATIC token. However, it lacks passive revenue-sharing mechanisms, such as royalties for derivative works.

Commercial startups have emerged aligned to Content ARCs to deliver AI contributor compensation. Bria (www.bria.ai) and ProRata (www.prorata.ai) state that they train their AI models on curated licensed content, and so the issue of Authenticity and content identification within their centralized systems is moot. For Compensation, both companies use proprietary attribution mechanisms to evaluate the relevance of training data to generated outputs. Bria distributes revenue through a scoring mechanism, compensating contributors at both training and inference time based on their data’s influence. ProRata provides credit and compensation on a per-use basis. However, neither enables granular rights management. Dataswyft (www.dataswyft.com) offers a contract-based data ecosystem facilitating self-sovereign data ownership without scalability concerns of DLT. For rights, the HAT Microserver Instruction Contracts define the terms under which external parties can access a user’s data to train AI but require manual enforcement and renewal by Dataswyft. The system facilitates compensated access to data within a contractually enforced marketplace. Other startups and image stock companies, such as Getty (www.gettyimages.co.uk/ai), have introduced their own generative AI services that claim to compensate creators whose data was used in training the models; however, due to their proprietary nature, there is a lack of public detail on the compensation model. In general, commercial platforms focus on either or both of the rights and compensation phases of Content ARCs, rather than provenance and authenticity of data across platforms.

5 Conclusion

Once implemented, Content ARCs would enable end-to-end machine-readable permissions and automatic compensation throughout the AI development. A future case study could involve a content creator utilizing ARC-compliant tools to manage rights and receive compensation for their work used in AI training. This demonstrates how our framework directly addresses current policy needs, responding to the needs outlined in the UK’s recent consultation on AI and copyright, which addresses transparent licensing and fair remuneration for creators in AI development [2].

No single open standard or system for content rights yet delivers fully across all phases and capabilities of the Content ARC framework. Several technical barriers remain, of which we list the most significant. First, the challenge of establishing registries for authenticity, ownership and licenses leads most solutions toward centralized or DLT solutions deployed on private chains due to scalability concerns. Yet this raises governance issues around liability and operation of the registry. Identity is key to representing rightsholders and licensees, yet self-sovereign identity (SSI) identity frame-

works remain under-adopted, including in the creative economy. Today, most decentralized systems are leveraging the blockchain wallet address as an identifier. Ongoing legislative debate around copyright form may hinder adoption and effectiveness of rights declarations, which are reliant upon legal remedy rather than technology (*i.e.* DRM) to ensure compliance. A further limitation is that content discovery mechanisms are largely absent today; even NFT markets were centralized search portals, indexing decentralized registries. Finally, despite a few early scoping workshops exploring value in decentralized licensing [27,32] there is a lack of ‘in the wild’ feasibility studies exploring how this would map to business models based on Content ARCs—answering this will be key to creating impact in the creative economy.

Acknowledgements

This work was supported by UKRI Grant EP/T022485/1.

References

- [1] M. M. Grynbaum and R. Mac, “The Times Sues OpenAI and Microsoft Over A.I. Use of Copyrighted Work,” *The New York Times*, 2023. [Online]. Available: <https://www.nytimes.com/2023/12/27/business/media/new-york-times-open-ai-microsoft-lawsuit.html>
- [2] UK Intellectual Property Office, “Consultation on Copyright and Artificial Intelligence,” <https://www.gov.uk/government/consultations/copyright-and-artificial-intelligence>, February 2025.
- [3] Financial Times, “AI copyright wars need a market solution,” *Financial Times*, 2023. [Online]. Available: <https://www.ft.com/content/304d660f-6cac-4e38-a6d5-d8d98f5770fb>
- [4] A. C. D. Giustina, “Fair Compensation for Copyrighted Data Used in AI Training,” Master’s Thesis, Tilburg University, 2024.
- [5] A. Cranz, “EU to invest €200 billion in AI development to compete with us and china,” *The Verge*, 2025. [Online]. Available: <https://www.theverge.com/news/609930/eu-200-billion-investment-ai-development>
- [6] T. Carpenter, “We could use a model licensing framework for scholarly content use in AI tools,” 2025, available at: <https://scholarlykitchen.sspnet.org/2025/02/26/we-could-use-a-model-licensing-framework-for-ai-tools/>.
- [7] W3C, “TDM Reservation Protocol (TDMRep),” <https://www.w3.org/community/reports/tdmrep/CG-FINAL-tdmrep-20240202>, February 2024.
- [8] “Directive (EU) 2019/790 of the European Parliament and of the Council of 17 April 2019 on Copyright and Related Rights in the Digital Single Market,” *Official Journal of the European Union*, 2019. [Online]. Available: <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX%3A32019L0790>
- [9] International Press Telecommunications Council (IPTC), *IPTC Photo Metadata Standard v2024.1*, 2024. [Online]. Available: <https://iptc.org/standards/photo-metadata/>
- [10] *Coalition for Content Provenance and Authenticity, Technical Specification v2.1*, <https://c2pa.org>, 2024.
- [11] T. Bui, S. Agarwal, and J. Collomosse, “TrustMark: Universal Watermarking for Arbitrary Resolution Images,” *arXiv preprint arXiv:2311.18297*, 2023. [Online]. Available: <https://arxiv.org/abs/2311.18297>
- [12] P. Fernandez, H. Elshar, I. Z. Yalniz *et al.*, “Video seal: Open and efficient video watermarking,” *arXiv preprint arXiv:2412.09492*, 2024. [Online]. Available: <https://arxiv.org/abs/2412.09492>
- [13] A. Black, T. Bui, H. Jin *et al.*, “Deep Image Comparator: Learning to Visualize Editorial Change,” in *Proc. CVPRW (Workshop on Media Forensics)*, 2021.
- [14] “International Standard Content Code (ISCC),” <https://iscc.codes/>.
- [15] E. Nguyen, T. Bui, V. Swaminathan *et al.*, “OSCAR-Net: Object-centric Scene Graph Attention for Image Attribution,” in *Proc. ICCV*, 2021.
- [16] K. Balan, A. Black, S. Jenni *et al.*, “DECORAIT - DECentralized Opt-in/out Registry for AI Training,” in *Proc. of the Conference on Visual Media Production (CVMP)*, 2023.
- [17] PLUS Coalition, “PLUS Coalition,” <https://www.useplus.com/>, 2020.
- [18] J. Collomosse and A. Parsons, “To Authenticity, and Beyond! Building Safe and Fair Generative AI upon the Three Pillars of Provenance,” *IEEE Computer Graphics and Applications (IEEE CG&A)*, 2024.
- [19] P. Rixhon, “An update on JPEG trust,” https://cawg.io/meeting-notes/_attachments/2025-01-21/jpeg-trust-presentation.pdf, January 2025.
- [20] National Security Agency, “Content credentials: Strengthening multimedia integrity in the generative ai era,” Tech. Rep., Jan. 2025. [Online]. Available: <https://media.defense.gov/2025/Jan/29/2003634788/-1/-1/0/CSI-CONTENT-CREDENTIALS.PDF>
- [21] C. Zauner, “Implementation and benchmarking of perceptual image hash functions,” Master’s thesis, Upper Austria University of Applied Sciences, Hagenberg, 2010, original pHash documentation and implementations available at <http://www.phash.org>. [Online]. Available: <https://phash.org/>
- [22] M. Steinebach, “An analysis of photodna,” in *Proc. of the 18th International Conference on Availability, Reliability and Security*, ser. ARES ’23. New York, NY, USA: Association for Computing Machinery, 2023.
- [23] Facebook Open Source, “PDQ: Perceptual Hash for Image Near-Duplicate Detection,” <https://github.com/facebook/ThreatExchange>, 2019.
- [24] H. Liu, R. Wang, S. Shan *et al.*, “Deep supervised hashing for fast image retrieval,” in *Proc. CVPR*, 2016.
- [25] Creative Commons, “Machine-readable metadata,” 2008. [Online]. Available: <https://creativecommons.org/about/downloads/>
- [26] R. Iannella and V. Rodríguez-Doncel, “ODRL vocabulary & expression 2.2,” W3C Recommendation, February 2018. [Online]. Available: <https://www.w3.org/ns/odrl/2/>
- [27] Digital Catapult, “DLT Field Lab Report: Emerging Futures for Tokenisation and Digital Media Rights,” UKRI DECade, Tech. Rep., 2024. [Online]. Available: www.digicatapult.org.uk/wp-content/uploads/2024/12/Catapult.ORAgen_Emerging_Futures_Report.pdf
- [28] K. Balan, S. Agarwal, S. Jenni *et al.*, “EKILA: Synthetic Media Provenance and Attribution for Generative Art,” in *Proc. CVPR Workshop on Media Forensics (CVPRW)*, 2023.
- [29] S. Wang, A. Efros, J. Zhu *et al.*, “Evaluating data attribution for text-to-image models,” in *Proc. ICCV*, 2023.
- [30] V. Asnani, J. Collomosse, T. Bui *et al.*, “ProMark: Proactive diffusion watermarking for causal attribution,” in *Proc. CVPR*, 2024.
- [31] J. T. Wang, Z. Deng, H. Chiba-Okabe *et al.*, “An economic solution to copyright challenges of generative ai,” *arXiv preprint arXiv:2404.13964*, 2024. [Online]. Available: <https://arxiv.org/abs/2404.13964>
- [32] F. Liddell, E. Tallyn, E. Morgan *et al.*, “ORAgan: Exploring the Design of Attribution through Media Tokenisation,” in *Proc. ACM Designing Interactive Systems (DIS)*, 2024.