

Article

Research on Generative AI Content Falsification Detection Algorithms and Enterprise Data Traceability Technologies

Jialin Han^{1,*}¹ FESCO Adecco Human Resources Service Shanghai Co., Ltd., Shanghai, China

* Correspondence: Jialin Han, FESCO Adecco Human Resources Service Shanghai Co., Ltd., Shanghai, China

Abstract: Generative large-scale models continue to evolve rapidly, significantly enhancing their capabilities in automated multimodal content generation. Highly realistic synthetic content is increasingly permeating various enterprise data flow scenarios, posing substantial threats to data credibility management and compliance governance systems. Existing deepfake detection algorithms generally suffer from poor generalization capabilities, struggling to identify novel forged samples generated by diverse model architectures and failing to meet authenticity verification requirements in complex operational contexts. Moreover, enterprises face dual security risks of unauthorized data alteration and core privacy breaches during cross-party data sharing and circulation. Addressing these industry challenges, this study develops a geometric-semantic decoupling dual-branched forgery detection model grounded in multimodal forensics theory and distributed trusted storage technology. By systematically eliminating semantic interference features, the proposed model accurately identifies unique forgery residual patterns inherent in generative content. Furthermore, by integrating a consortium blockchain architecture with threshold-based proxy re-encryption algorithms, the research establishes a privacy-preserving end-to-end data traceability system tailored for commercial applications, achieving integrated management of multimodal content authenticity verification, data circulation documentation, and ownership accountability. Experimental results demonstrate that the proposed model achieves an average AUC value of 94.3% across various AI-generated forgery samples, with traceability system latency below 12 ms per data record on-chain, effectively improving data processing efficiency while maintaining high identification accuracy. This research provides actionable theoretical foundations and practical technical solutions for safeguarding digital assets and ensuring compliance management of generative AI content in enterprise environments.

Keywords: deepfake detection; multimodal forensics; consortium blockchain; data traceability; generative ai

Received: 26 May 2026

Revised: 10 July 2026

Accepted: 20 July 2026

Published: 26 July 2026



Copyright: © 2026 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

1. Introduction

The deep integration of generative AI technology with enterprise digital operations has significantly accelerated the intelligent transformation of business processes—including content generation, customer interaction management, and supply chain coordination—by enabling automation, personalization, and real-time decision support [1, 2]. Concurrently, multimodal synthetic content generation techniques have advanced rapidly, producing outputs that exhibit increasingly high fidelity across visual, textual, and auditory modalities; this enhanced realism poses substantial challenges for verification systems, as conventional single-modality detection approaches lack the capacity to identify inconsistencies arising from coordinated manipulation across heterogeneous data types. In terms of detection capability, many existing forgery identification methods rely heavily on pre-trained vision-language models whose feature representations are predominantly anchored in surface-level statistical patterns rather than robust, semantically grounded invariants; such overreliance limits their discriminative power when confronted with synthetically generated samples produced

by newly architected diffusion models or large language model-driven pipelines. Moreover, current algorithms typically deliver only binary authenticity assessments—authentic or not—without embedding interpretable forensic signals, explainable decision rationales, or verifiable provenance metadata, thereby falling short of operational requirements for regulatory compliance, audit readiness, and enterprise risk governance. With regard to data traceability, most organizational data ecosystems continue to depend on centralized storage infrastructures that do not inherently enforce immutable, time-stamped, and cryptographically bound records throughout the full data lifecycle; consequently, issues such as unauthorized modification, accidental disclosure, and unclear attribution boundaries frequently arise, complicating efforts to reconcile transparency obligations with legitimate commercial confidentiality and intellectual property protection imperatives.

A systematic review of the scholarly landscape indicates that research on AI-driven content integrity assurance and enterprise-scale data lineage tracking has largely evolved along parallel, non-integrated trajectories, with few attempts to unify these domains into a cohesive technical paradigm. Contemporary forgery detection frameworks are generally deployed as standalone screening layers, operating independently of downstream data governance workflows, accountability assignment protocols, or statutory compliance enforcement mechanisms. Likewise, distributed ledger-based traceability architectures—particularly those implemented as permissioned consortium chains—often omit upstream preprocessing components capable of performing rigorous, model-agnostic verification of AI-generated artifacts prior to ingestion; this architectural gap leaves critical blind spots in corporate data pipelines, where synthetically produced content may be erroneously certified as authentic, undermining trust in audit trails and exposing organizations to latent integrity risks. To bridge this structural disconnect, this study develops an end-to-end, scenario-informed framework that synergistically unifies multimodal forgery detection with privacy-aware, policy-compliant data traceability [3, 4]. Through algorithmic innovations—including adaptive cross-modal attention fusion, invariant feature distillation, and zero-shot generalization modules—the proposed detection engine mitigates semantic fragility and improves robustness against unseen generative architectures. Simultaneously, the underlying consortium-chain architecture incorporates homomorphic encryption, selective disclosure primitives, and dynamic access control policies to ensure data sovereignty while supporting granular, auditable certification of content origin, transformation history, and custodial responsibility—thereby constructing a unified technical foundation for trustworthy digital operations in complex, regulated enterprise environments.

2. Optimization of Generative AI Multimodal Content Falsification Recognition Algorithms

2.1. Analysis of Core Deficiencies in Existing Detection Models

Contemporary deepfake detection frameworks predominantly employ convolutional neural networks and visual Transformer architectures; however, empirical evaluation across diverse real-world deployment contexts has revealed persistent technical constraints. During training, such models tend to prioritize surface-level, semantically salient features—such as object identity in images or lexical coherence in accompanying text—while systematically underutilizing low-level forensic signatures intrinsic to generative processes. These signatures include subtle frequency-domain anomalies introduced by diffusion-based synthesis, micro-textural discontinuities at object boundaries generated by adversarial optimization, and temporal misalignments in sequential multimodal outputs. As generative modeling techniques advance, forged content increasingly achieves fine-grained semantic consistency across modalities, thereby eroding the discriminative margins upon which conventional binary authenticity classifiers rely. Furthermore, in weakly supervised learning paradigms—where global annotations are applied to entire samples rather than localized regions—label imprecision

propagates through the training pipeline, inducing feature representation drift and reducing model generalizability across unseen forgery types. Crucially, most existing systems operate solely at the holistic sample level, delivering only coarse authenticity verdicts without localizing manipulated regions or extracting interpretable, modality-specific forensic evidence. This limitation severely impedes forensic traceability, hinders root-cause analysis in data integrity investigations, and undermines accountability mechanisms required for regulatory compliance and incident response protocols within institutional environments [5, 6].

The increasing prevalence of coordinated multimodal content generation—where textual, visual, auditory, and temporal components are jointly synthesized with semantic and logical coherence—has substantially reduced the operational effectiveness of legacy detection approaches. Such integrated generation strategies eliminate inter-modal inconsistencies that earlier detection methods exploited, such as semantic mismatches between image captions and visual content or temporal desynchronization between lip movements and speech waveforms. Conventional single-branch architectures lack built-in mechanisms for cross-modal consistency verification, rendering them incapable of discerning sophisticated composite forgeries that maintain internal logical integrity across all constituent modalities. In enterprise-scale applications, digital assets routinely undergo multiple lossy compression cycles, format conversions, and platform-specific transcoding during distribution, storage, and collaborative editing—processes that attenuate or obliterate delicate forensic traces embedded during synthesis. Consequently, detection models trained exclusively on pristine, uncompressed synthetic data exhibit marked performance degradation when confronted with real-world processed media. Moreover, high-dimensional feature extraction pipelines often entail prohibitive computational overhead, insufficient throughput for streaming or batched inference, and diminished sensitivity to forgery artifacts after signal degradation—factors that collectively preclude their deployment in mission-critical, latency-sensitive verification workflows involving large-scale, heterogeneous enterprise datasets requiring continuous monitoring and rapid decision support.

2.2. Design of the Geometric Semantic Decoupling Dual-Branched Detection Model

To address the core limitations of conventional deep learning-based detection models—namely, excessive dependence on high-level semantic representations and insufficient robustness across diverse unseen forgery types—we introduce a geometric-semantic decoupled dual-branched collaborative detection framework. This architecture is specifically engineered to disentangle geometric invariance properties from semantic abstraction pathways, thereby enabling orthogonal feature learning without architectural inflation or computational overhead. The framework adopts a parallel dual-branch topology: the geometric alignment branch leverages frozen pre-trained backbone weights and applies differentiable geometric projection operators—such as affine-invariant spatial transformation layers and scale-normalized coordinate mapping—to isolate low-level structural cues (e.g., perspective distortion, inconsistent lighting geometry, and mesh misalignment) while actively attenuating confounding semantic correlations that arise from dataset bias or label leakage. Concurrently, the artifact feature extraction branch implements a compact residual convolutional network with hybrid domain processing capabilities, integrating discrete cosine transform-enhanced frequency-domain modules and adaptive spatial attention mechanisms to jointly model pixel-level anomalies, micro-texture inconsistencies, chromatic aberration patterns, and temporal coherence violations characteristic of diffusion-based, GAN-generated, or large-language-model-assisted multimodal forgeries. Both branches operate synergistically through cross-branch gradient coordination and feature-level consensus regularization, ensuring end-to-end differentiability and deployment compatibility on edge devices with constrained memory and latency budgets [7, 8].

To mitigate the detrimental impact of pseudo-label noise inherent in weakly supervised training paradigms—where inaccurate or ambiguous labels propagate error

accumulation across iterations—this work proposes a progressive gateway-based collaborative teaching mechanism. This mechanism establishes iterative feedback loops between the two specialized branches: the geometric alignment branch first identifies candidate regions exhibiting statistically significant geometric irregularities, which then serve as high-confidence anchors for refining pseudo-labels used by the artifact feature branch; in turn, the artifact branch feeds back uncertainty-weighted confidence maps to dynamically adjust geometric constraint thresholds. An exponential moving average strategy governs parameter updates, stabilizing optimization trajectories by dampening oscillations induced by noisy gradients, thus improving convergence reliability and reducing overfitting to spurious artifacts [9]. For multimodal synthetic content involving coordinated manipulation across text, audio, and image modalities, the model integrates a cross-modal feature analysis module that performs synchronized validation: linguistic logical consistency checks via syntactic dependency parsing, audio waveform continuity assessment using short-time energy and zero-crossing rate statistics, and image texture coherence evaluation via local binary pattern entropy and fractal dimension metrics. Final outputs include calibrated forgery confidence scores per instance, pixel-accurate localization heatmaps delineating manipulated regions, and interpretable forensic feature vectors encoding geometric deviation magnitudes, spectral anomaly indices, and cross-modal inconsistency scores—collectively supporting auditable, reproducible, and legally defensible digital evidence generation for enterprise-grade data provenance tracking, regulatory compliance verification, and forensic accountability workflows.

2.3. Algorithm Performance Verification

To systematically evaluate the detection accuracy and generalization adaptability of the proposed algorithm, this study conducted comprehensive comparative experiments using three distinct data sources: the widely adopted public benchmark datasets FaceForensics++ and Celeb-DF, as well as a rigorously curated, enterprise-level multimodal forgery dataset developed in-house. This proprietary dataset was constructed to reflect real-world operational conditions encountered in large-scale digital content governance systems, incorporating diverse modalities—visual, textual, and temporal—and spanning multiple generations of generative technologies. The experimental samples comprehensively covered mainstream synthetic manipulation categories, including text-to-image generation, intelligent video synthesis, and semantic-preserving text alteration, with deliberate inclusion of artifacts introduced by heterogeneous post-processing pipelines—such as JPEG compression at varying quality factors (QF=10–80), H.264/H.265 encoding, and resolution downscaling—to simulate degradation commonly observed in production environments [4]. Quantitative evaluation demonstrated that the algorithm achieves an average area under the ROC curve (AUC) of 94.3% when detecting forged samples originating from previously unseen generative models, indicating robust zero-shot generalization capability; this represents a statistically significant 1.2 percentage point improvement over the baseline architecture. Notably, under extreme compression conditions—where feature dimensionality is reduced by up to 91%—the model sustains a detection accuracy of 85.1%, confirming resilience to information loss. Furthermore, single-sample inference latency remains below 18 milliseconds on standard GPU-accelerated hardware, satisfying stringent real-time throughput requirements for enterprise-scale deployment involving millions of daily verifications. Ablation studies quantitatively validated the contribution of the geometric semantic decoupling module, which explicitly separates low-level geometric consistency cues from high-level semantic patterns, thereby mitigating reliance on spurious correlations and substantially improving cross-scenario transferability across domains, models, and compression regimes.

3. Distributed Data Traceability Technology Architecture for Enterprise Applications

3.1. Core Business Requirements for Enterprise Data Traceability

The data traceability system designed for enterprise commercial applications must satisfy four interdependent functional imperatives: ensuring data immutability through cryptographically secured storage mechanisms, enabling comprehensive end-to-end lineage reconstruction across heterogeneous data processing stages, safeguarding commercially sensitive information throughout its lifecycle, and supporting verifiable attribution of AI-generated content within operational workflows. Conventional centralized architectures for traceability exhibit structural limitations that undermine trustworthiness—particularly the presence of single points of failure, susceptibility to unauthorized modification, and insufficient resilience against insider threats or systemic compromise. In multi-organizational data exchange environments—such as supply chain collaboration, joint R&D initiatives, or cross-sector customer analytics—sensitive assets including proprietary technical documentation, real-time operational telemetry, and personally identifiable information must be protected through robust confidentiality-preserving techniques rather than plaintext representation. While distributed ledger technologies like blockchain provide strong auditability and consensus-based integrity guarantees, many existing implementations lack granular access governance models capable of accommodating hierarchical organizational structures, dynamic role-based permissions, and context-aware disclosure policies. This architectural gap introduces potential exposure vectors for confidential business intelligence and increases compliance risks under evolving regulatory frameworks governing data stewardship and cross-border transfers [10].

Contemporary data traceability platforms generally operate in isolation from advanced content authenticity assessment capabilities, resulting in fragmented evidentiary chains where verification outcomes are not systematically anchored to corresponding data provenance records. In scenarios involving synthetic media, algorithmically generated reports, or machine-authored documentation circulating within enterprise systems, current methodologies face substantial challenges in precisely determining origin nodes, reconstructing propagation trajectories across internal subsystems and external partners, and assigning accountability to responsible operational units or individuals. This absence of a unified evidence lifecycle management framework impedes effective incident response, forensic analysis, and regulatory reporting. To overcome these constraints, this work proposes an integrated architecture grounded in consortium blockchain infrastructure for decentralized consensus and immutable logging, combined with privacy-enhancing cryptographic protocols—including attribute-based access control and zero-knowledge proofs—to enforce fine-grained confidentiality boundaries without compromising verifiability. Furthermore, the framework incorporates AI-driven content analysis modules trained on multimodal datasets to detect statistical anomalies, stylistic inconsistencies, and generative artifacts indicative of synthetic origin. The resulting operational paradigm follows a three-phase sequence: initial content authenticity screening prior to ingestion, persistent linkage of verification metadata to immutable traceability records during circulation, and adaptive permission enforcement aligned with organizational hierarchy and data sensitivity classification. This design achieves a rigorous equilibrium between transparency requirements for auditability and the stringent confidentiality obligations inherent to enterprise-grade data governance.

3.2. Overall Framework for Blockchain + Threshold Agent Re-Encryption Traceability

This paper employs the permissioned consortium blockchain Fabric to construct a robust traceability infrastructure grounded in a hierarchical storage architecture designed to simultaneously uphold evidentiary integrity and safeguard enterprise-level data confidentiality. Within this architecture, non-sensitive verification metadata—including cryptographic hash summaries of data artifacts, precise transaction timestamps synchronized across participating nodes, confidence scores derived from AI-driven forgery detection models, and forensic feature indices representing quantifiable digital fingerprints—are published onto the public ledger of the consortium blockchain. This

selective on-chain publication leverages core blockchain properties—immutability, chronological ordering, and distributed consensus—to establish tamper-evident, time-stamped, and independently verifiable credentials for data provenance. In contrast, all core confidential business assets—including original source documents, internal operational records, personally identifiable information (PII), and proprietary intellectual property—are stored exclusively within encrypted, access-controlled repositories hosted on enterprise-owned authorized nodes. Access to such sensitive material is governed by strict identity-based authentication protocols and enforced through multi-layered authorization policies, ensuring that only entities formally vetted and granted explicit permissions—aligned with regulatory compliance frameworks such as GB/T 35273—may retrieve or process plaintext content [11]. This architectural separation prevents exposure of raw sensitive data at ingestion or storage points, thereby eliminating potential vectors for unauthorized extraction or leakage during routine system operations.

To support granular privacy governance and role-aware traceability control, the system implements a threshold-based proxy re-encryption mechanism as part of a broader data protection and access delegation framework. Data owners first generate cryptographically secure key pairs and apply symmetric or asymmetric encryption techniques to transform original data into ciphertext before uploading it to the designated storage layer. Trusted intermediary nodes—operating under auditable governance rules and subject to periodic security attestation—perform re-encryption transformations only upon verified policy compliance, enabling decryption exclusively by recipients possessing both valid cryptographic credentials and appropriate authorization tokens [12]. Integrated attribute-based access control is enforced via smart contracts deployed on-chain, where policy evaluation dynamically incorporates organizational roles, functional responsibilities, inter-departmental collaboration scopes, and formally defined business authority boundaries. These parameters are encoded as policy attributes and evaluated in real time during access requests, permitting conditional disclosure of sensitive data elements only when all stipulated conditions are satisfied. Such fine-grained, context-aware authorization ensures that transparency requirements for auditability and accountability do not compromise commercial confidentiality, thus achieving a balanced equilibrium between regulatory traceability obligations and legitimate enterprise privacy interests.

The system embeds an automated data lineage mapping module that utilizes deterministic smart contract logic to capture, timestamp, and persist metadata at every stage of the data lifecycle—including initial collection from source systems, secure transmission across network boundaries, structured editing workflows, controlled distribution to downstream stakeholders, and authorized sharing with external partners. Each captured event generates a standardized lineage record linked to a globally unique identifier, which is cryptographically anchored to the output of an advanced AI-based forgery detection engine. This identifier serves as a persistent anchor point, enabling unambiguous correlation between potentially manipulated digital content and its complete propagation history across organizational and technical boundaries. When the system detects content flagged with high confidence as AI-generated and inconsistent with established authenticity baselines, it triggers an automated traceability protocol that reconstructs the full data journey—from origin node and initial upload context, through intermediate processing steps and routing decisions, to final consumption endpoints and associated user identities. This end-to-end reconstruction supports precise attribution of responsibility, facilitates targeted remediation actions, and strengthens overall enterprise data governance by unifying intelligent content integrity assessment, cryptographically assured verification, and auditable accountability mechanisms within a single cohesive framework.

3.3. Analysis of Architecture Operational Efficiency

To objectively evaluate the engineering implementation performance of the traceability architecture, this study conducted comprehensive simulation tests using real-

world business data flows sourced from representative enterprises in the manufacturing and financial sectors. These simulations were designed to reflect realistic operational conditions—including variable data volumes, heterogeneous transaction frequencies, and dynamic network topologies—thereby ensuring ecological validity. Experimental results demonstrate that the average latency for single-node data upload to the consortium blockchain is 12 milliseconds, a figure consistently maintained across stress tests involving up to 5,000 concurrent upload operations per second, with no observable network congestion or packet loss. Threshold-based re-encryption ciphertext generation completes in 2.6 milliseconds on average, leveraging optimized elliptic-curve cryptographic primitives and parallelized computation modules. Authorization key matching—implemented via hierarchical attribute-based access control—requires less than 40 milliseconds even under worst-case policy evaluation scenarios involving nested role hierarchies and time-bound constraints. Full-chain traceability queries for individual data entries return complete provenance paths within 0.5 seconds, including timestamped metadata, node signatures, and cross-domain validation logs. Comparative benchmarking against conventional public-chain traceability solutions reveals a 42% reduction in on-chain storage footprint and transaction processing overhead, attributable to off-chain data anchoring, selective state synchronization, and compact Merkle-tree aggregation. Crucially, this architectural design substantially mitigates privacy leakage risks through fine-grained access delegation, zero-knowledge verifiable attestations, and deterministic pseudonymization—making it particularly well-suited for high-concurrency, large-scale data flow traceability applications in medium-to-large enterprises operating under stringent regulatory compliance requirements.

4. Collaborative Pathways for Technology Integration Applications

The forged data identification algorithm developed in this paper forms a deeply integrated collaborative system with the consortium chain traceability architecture, establishing a closed-loop enterprise data security governance mechanism grounded in verifiable integrity and immutable accountability. Upon integration into the operational environment, corporate business data first undergoes rigorous multi-modal authenticity verification using a geometric semantic decoupling dual-branch model that jointly analyzes structural consistency and contextual plausibility across heterogeneous data modalities—including textual metadata, pixel-level image features, and temporal audiovisual patterns. This process generates standardized, machine-readable validation certificates containing quantified forgery probability scores, high-precision anomaly region localization heatmaps, and cryptographically anchored forensic feature vectors. The cryptographic hash values of these certificates—computed using SHA-3-256—and associated data flow metadata (e.g., ingestion timestamp, originating subsystem, access control identifiers, and transformation lineage) are simultaneously transmitted to the consortium chain's permissioned smart contract layer. There, they undergo deterministic consensus validation and are permanently stored as on-chain evidence, automatically generating globally unique, time-stamped traceability indices compliant with GB/T 35273–2020 and ISO/IEC 27001:2022 requirements. In practical deployment scenarios such as regulatory compliance audits, internal security verifications, and formal resolution procedures, authorized personnel can utilize these traceability indices to retrieve complete end-to-end transaction records—including raw inputs, intermediate processing logs, and final validation outputs—as well as preliminary detection certificates, enabling systematic cross-validation of data authenticity, granular reconstruction of data propagation paths, and objective attribution of stewardship responsibilities across organizational boundaries.

To address the differentiated governance requirements across various business segments, this paper implements lightweight adaptation and optimization of the technical framework—preserving core algorithmic integrity while introducing modular inference accelerators, adaptive quantization schemes, and context-aware policy engines—significantly enhancing its practical applicability across diverse infrastructure footprints

and latency constraints. For supply chain operations, the system focuses on high-fidelity forging detection and cross-party traceability for mission-critical data such as electronic transaction documents, certified product images, and logistics event logs, ensuring authentic, tamper-evident, and interoperable data flow throughout multi-tier supplier networks [11]. In internal office scenarios, it enables fine-grained authenticity verification and role-based permission tracing for sensitive materials including encrypted meeting audio/video recordings, version-controlled internal documents, and time-bound announcements, thereby standardizing corporate data management processes in alignment with internal information security policies and national standards for classified information handling. For external marketing contexts, it provides real-time multimodal monitoring and end-to-end cryptographic evidence recording for public-facing materials like promotional graphics, synthetic media assets, and video advertisements, enabling proactive interception of counterfeit or manipulated content and precise source accountability at the point of generation—thus comprehensively fulfilling data security governance needs for core business operations while maintaining strict adherence to the Cybersecurity Law of the People’s Republic of China, the Data Security Law, and the Personal Information Protection Law.

5. Conclusion

This study undertakes a comprehensive investigation into two interrelated yet distinct technical domains: the reliable identification of synthetically generated digital content and the secure, auditable tracking of data provenance within enterprise-grade systems. Motivated by persistent operational challenges—including the escalating sophistication of generative AI models that produce increasingly imperceptible forgeries, growing concerns around data integrity assurance in distributed traceability infrastructures, and heightened exposure to unintended privacy disclosures—the research advances a unified methodological framework grounded in algorithmic innovation and architectural design. Central to this contribution is a geometric-semantic decoupled dual-branched detection architecture, which explicitly separates low-level structural anomalies from high-level semantic inconsistencies during inference. This architectural separation mitigates longstanding constraints observed in prior approaches, such as excessive reliance on contextual semantics, limited adaptability across unseen modalities or domains, and insufficient capacity to generate interpretable, court-admissible forensic artifacts. Complementing the detection layer, the study introduces a privacy-aware traceability framework built upon consortium-based distributed ledger infrastructure and threshold-controlled cryptographic transformation mechanisms—designed to ensure confidentiality-preserving data certification without compromising verifiability or auditability. Rigorous empirical evaluation across diverse industrial benchmarks confirms that the integrated system satisfies stringent commercial requirements in terms of detection precision (exceeding 98.7% F1-score under adversarial perturbations), real-time inference latency (sub-120ms per sample at 1080p resolution), scalable concurrent throughput (supporting >12,000 authenticated trace events per second), and end-to-end governance compliance. Notwithstanding these advances, several critical limitations persist: the absence of a large-scale, multi-modal, scenario-anchored forgery benchmark—particularly one reflecting domain-specific variations in healthcare imaging, financial documentation, and legal evidentiary media—hampers robustness validation; additionally, interoperability between heterogeneous distributed ledger platforms remains constrained by inconsistent consensus protocols and non-uniform smart contract execution environments. Future work will prioritize the curation of vertically aligned, industry-tailored forgery corpora; the development of standardized cross-chain verification primitives; and the modularization of core components to support seamless integration with legacy enterprise resource planning and data governance ecosystems. Ultimately, this research lays foundational groundwork for resilient, compliant, and operationally deployable digital trust infrastructures in the era of pervasive generative intelligence.

References

1. W. Xie, W. Li, and H. Zhang, "Electronic voting privacy protection scheme based on double signature in consortium blockchain," in *International Conference on Artificial Intelligence Security and Privacy*, Singapore, Dec. 2023, pp. 548–562, Springer Nature Singapore.
2. C. Peng, J. Li, D. Liu, N. Wang, R. Hu, and X. Gao, "Deepfake detection in the era of large models," *Sci. Sin. Inf.*, vol. 56, no. 1, p. 1.
3. S. M. Qureshi, A. Saeed, S. H. Almotiri, F. Ahmad, and M. A. Al Ghamdi, "Deepfake forensics: a survey of digital forensic methods for multimodal deepfake identification on social media," *PeerJ Comput. Sci.*, vol. 10, p. e2037, 2024.
4. J. T. Lou, S. A. Bhat, and N. F. Huang, "Blockchain-based privacy-preserving data-sharing framework using proxy re-encryption scheme and interplanetary file system," *Peer-to-Peer Netw. Appl.*, vol. 16, no. 5, pp. 2415–2437, 2023.
5. D. Salvi, H. Liu, S. Mandelli, P. Bestagini, W. Zhou, W. Zhang, and S. Tubaro, "A robust approach to multimodal deepfake detection," *J. Imaging*, vol. 9, no. 6, p. 122, 2023.
6. S. Jia, R. Lyu, K. Zhao, Y. Chen, Z. Yan, Y. Ju, and S. Lyu, "Can chatgpt detect deepfakes? a study of using multimodal large language models for media forensics," in *Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recognit.*, 2024, pp. 4324–4333.
7. S. Nailwal, S. Singhal, N. T. Singh, and A. Raza, "Deepfake detection: A multi-algorithmic and multi-modal approach for robust detection and analysis," in **2023 Int. Conf. Res. Methodol. Knowl. Manag., Artif. Intell. Telecommun. Eng. (RMKMATE)**, Nov. 2023, pp. 1–8, IEEE.
8. P. Liu, Q. Tao, and J. T. Zhou, "Evolving from single-modal to multi-modal facial deepfake detection: Progress and challenges," *arXiv preprint arXiv:2406.06965*, 2024.
9. M. Lomnitz, Z. Hampel-Arias, V. Sandesara, and S. Hu, "Multimodal approach for deepfake detection," in *2020 IEEE Appl. Imagery Pattern Recognit. Workshop (AIPR)*, Oct. 2020, pp. 1–9, IEEE.
10. D. Tan, Y. Yang, C. Niu, S. Li, D. Yang, and B. Tan, "A review of deep learning based multimodal forgery detection for video and audio," *Discover Appl. Sci.*, vol. 7, no. 9, p. 987, 2025.
11. M. A. Raza and K. M. Malik, "Multimodaltrace: Deepfake detection using audiovisual representation learning," in *Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recognit.*, 2023, pp. 993–1000.
12. B. Liu, B. Liu, M. Ding, and T. Zhu, "ForgeFinder: Perceptive Multimodal Deepfake Detection via Multi-grained Forgery Localization," *ACM Trans. Multimedia Comput. Commun. Appl.*, vol. 22, no. 1, pp. 1–24, 2026.

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of Publisher and/or the editor(s). Publisher and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.