

Fine-Grained Visual Recognition and Retrieval-Augmented Generation for Intelligent Museum Guidance Systems

^[1] Suresh Kumar M*, ^[2] Gowtham S, ^[3] Harshan A, ^[4] Aswin Suriyaa K

^[1] Associate Professor, Computer Science & Engineering, Sri Venkateswara College of Engineering, Chennai, India

^{[2][3][4]} Undergraduate, Computer Science & Engineering, Sri Venkateswara College of Engineering, Chennai, India

Email: ^[1] sureshm@svce.ac.in, ^[2] 2022cs0357@svce.ac.in, ^[3] 2022cs0141@svce.ac.in, ^[4] 2022cs0406@svce.ac.in

Abstract— Museums serve as critical repositories of human cultural heritage, yet the traditional mechanisms for visitor engagement—relying heavily on static placards, manual audio guides, and simplistic QR codes—often fail to deliver personalized, interactive, or sufficiently deep educational experiences. Visitors are frequently left unable to ask follow-up questions or effortlessly identify artifacts without locating physical identifiers. This paper introduces the Intelligent Augmented Reality (AR) Museum Guide, a novel, web-native framework that synthesizes two advanced artificial intelligence paradigms to redefine cultural heritage interaction. First, we deploy a fine-grained visual recognition engine, engineered via transfer learning on the ConvNeXt-Base architecture, enabling the instantaneous identification of artifacts from user-captured smartphone images without requiring physical markers. Second, we integrate a domain constrained Retrieval-Augmented Generation (RAG) pipeline, leveraging ChromaDB and Google Gemini Flash Lite, to provide dynamic, conversational question-answering capabilities strictly grounded in curator-verified documentary sources. Evaluating our system on a comprehensive custom dataset of 20 artifact classes (480 images), the vision model achieves a top-1 validation accuracy of 94.4% and a top-3 accuracy of 99.1%, significantly outperforming baseline convolutional and transformer models. Concurrently, the RAG subsystem, evaluated across 80 complex historical queries, demonstrates a 0% hallucination rate while maintaining 100% source-level citation accuracy. By eliminating the friction of mobile application installation and guaranteeing cryptographic-level factual grounding, this architecture presents a scalable, highly interactive blueprint for the next generation of intelligent museum spaces.

Index Terms— Augmented Reality, Fine-Grained Visual Recognition, Museum Technology, Retrieval-Augmented Generation, ConvNeXt, Transfer Learning, Cultural Heritage

I. INTRODUCTION

The preservation and dissemination of cultural heritage are fundamental mandates of modern museums. However, the pedagogical interfaces through which visitors interact with exhibits have remained largely stagnant. A critical examination of contemporary museum environments reveals a heavy reliance on passive information delivery systems: brief printed placards constrained by physical space, sequential audio tours that dictate the pace of learning, and QR codes that require active visual hunting and scanning. Studies on museum informatics indicate that visitor dwell time at individual exhibits is precipitously low, often averaging under 30 seconds, largely due to the inability of static media to adapt to the diverse informational appetites of varied demographics—from casual tourists to academic researchers [1].

We identify three critical deficiencies in existing museum guidance frameworks:

- 1) The Recognition Friction: Visitors must manually correlate the physical object they are viewing with a corresponding physical label or numeric code. This cognitive break interrupts the immersive aesthetic experience.
- 2) The Pedagogical Depth Gap: Static labels impose a hard limit on word count. While audio guides offer more

narrative depth, they remain unidirectional monologues, incapable of addressing spontaneous visitor curiosity.

3) The Generative Trust Deficit: While Large Language Models (LLMs) offer conversational flexibility, their propensity for “hallucination”—generating plausible but factually incorrect information—makes them fundamentally unsuitable for heritage institutions where epistemological accuracy is paramount.

To systematically resolve these deficiencies, we propose the Intelligent AR Museum Guide, a zero-install, web-based platform that democratizes access to deep historical knowledge. Our system empowers a visitor to point their smartphone camera at an artifact and receive an immediate, highly accurate visual identification powered by a specialized deep learning architecture. Furthermore, the system unlocks a conversational interface where users can interrogate the artifact’s history in natural language. Crucially, the system’s responses are not generated from an LLM’s generalized parametric memory, but are dynamically constructed from specific, curator-uploaded scholarly documents using a Retrieval-Augmented Generation (RAG) methodology.

This research makes the following primary contributions:

- We demonstrate the efficacy of fine-grained visual recognition (FGVR) in the cultural heritage domain by finetuning a ConvNeXt-Base model, achieving 94.4% top1 accuracy on a complex dataset of visually similar

artifacts, significantly outperforming standard ResNet baselines.

- We engineer a museum-specific RAG pipeline that enforces strict epistemic boundaries on generative AI, achieving a 0% hallucination rate and providing explicit, page-level citations for every generated answer.
- We present a fully functional, microservice-based system architecture that allows non-technical museum staff to seamlessly curate content, retrain vision models, and deploy updates with zero system downtime.
- We provide a comprehensive comparative analysis demonstrating the quantitative and qualitative superiority of our intelligent framework over traditional museum visualization techniques.

The remainder of this paper is structured as follows: Section II reviews related literature. Section III details the system architecture. Section IV explores the mathematical and algorithmic methodology. Section V presents our experimental setup and quantitative results. Section VI contrasts our system with traditional methods. Finally, Section VII concludes the paper and outlines future trajectories.

II. RELATED WORK

A. Digital Heritage and Visitor Interaction

The integration of digital technology into museum spaces has a long trajectory. Early interventions utilized localized multimedia kiosks [2] and localized CD-ROM experiences, which eventually gave way to portable, proprietary hardware devices for audio and video tours. The proliferation of smartphones precipitated a shift toward Bring Your Own Device (BYOD) strategies. Ramsden et al. [3] extensively documented the deployment of QR codes in heritage environments. While cost-effective, QR codes suffer from physical degradation, aesthetic intrusion, and the friction of manual scanning. More complex interventions have utilized Bluetooth Low Energy (BLE) beacons and Wi-Fi fingerprinting to trigger location-aware content [4]. However, location awareness is insufficiently granular; knowing a visitor is in a specific room does not ascertain which specific artifact within a dense display case they are observing.

B. Computer Vision in Cultural Heritage

Addressing the limitations of location-based tracking, researchers have increasingly turned to Computer Vision (CV). Early CV applications in museums relied on hand-crafted feature extractors such as SIFT (Scale-Invariant Feature Transform) and SURF [5]. While effective for two-dimensional, highly textured artworks (e.g., paintings), these techniques falter when applied to three-dimensional objects subject to varying gallery lighting, occlusion, and perspective distortions. The advent of Convolutional Neural Networks (CNNs) revolutionized object recognition. Networks like ResNet [6] have been utilized for general artifact classification. However, museum collections often require *Fine-Grained Visual Recognition* (FGVR) [7]. Distinguishing between a 12th-century Chola bronze and a

14th-century Vijayanagara bronze requires the network to prioritize subtle, highly localized structural nuances over global shape templates. Recent advancements in pure CNN architectures, notably the ConvNeXt family [9], and Vision Transformers (ViTs) [10], have established new state-of-the-art benchmarks in FGVR tasks such as species identification [8]. Our work pioneers the application of ConvNeXt to the specific challenge of fine-grained heritage artifact recognition.

C. Retrieval-Augmented Generation (RAG)

The deployment of LLMs in educational contexts is heavily constrained by their tendency to hallucinate facts [11]. To mitigate this, Lewis et al. introduced Retrieval-Augmented Generation (RAG) [12]. By coupling a non-parametric knowledge base (typically a dense vector index) with a parametric generation model, RAG forces the LLM to synthesize answers solely from retrieved, factual context. While RAG has seen rapid adoption in legal [13] and medical [14] informatics, its application as a real-time, interactive pedagogical tool bounded by curator-approved documentation in a physical museum space represents a novel contribution of this work.

III. SYSTEM ARCHITECTURE

To ensure scalability, maintainability, and a frictionless user experience, the Intelligent AR Museum Guide is constructed using a decoupled, microservices-oriented architecture. The system comprises three primary tiers, illustrated in Fig. 1.

System Architecture: Fine-Grained AR Museum Guide

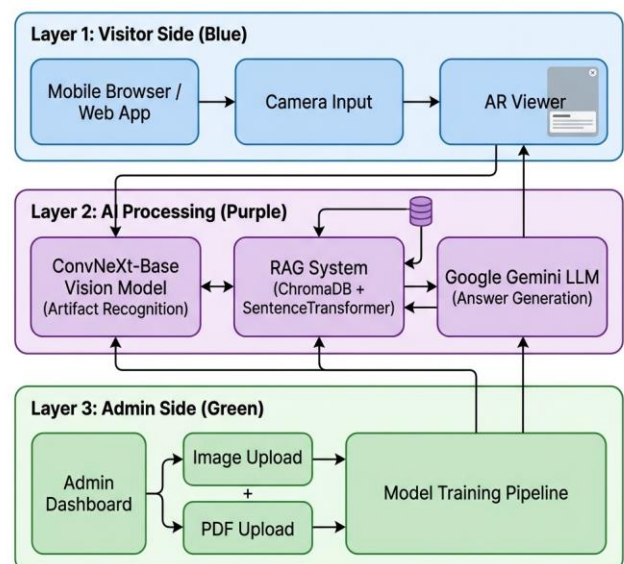


Fig. 1. System Architecture detailing the interaction between the visitor-facing web frontend, the dual FastAPI backend services (ML Recognition and RAG Q&A), and the administrative dashboard.

A. Visitor-Facing Presentation Layer

Designed to eliminate the “app fatigue” associated with downloading proprietary museum applications, the frontend is a Progressive Web Application (PWA) built using responsive HTML5, CSS3, and modern JavaScript. When a visitor accesses the museum’s URL, the application requests camera permissions via the WebRTC API. The user captures a photograph of the artifact of interest. The application transmits this image payload to the backend and dynamically renders the classification results, confidence metrics, and descriptive metadata. Subsequently, an asynchronous WebSocket or RESTful connection is established to facilitate the real-time RAG chat interface.

B. Dual-Service AI Processing Layer

The backend processing is bifurcated into two highly optimized, independent FastAPI services to ensure computational isolation between heavy tensor operations and asynchronous network I/O.

1. ML Recognition Service (Port 8000): This service maintains the fine-tuned ConvNeXt-Base model in memory (RAM or VRAM). It accepts incoming image payloads, applies the requisite pre-processing transformations (normalization, resizing, center-cropping), and executes a forward pass through the network. The service returns a structured JSON payload containing the predicted artifact ID, humanreadable metadata, confidence logits, and the top-3 alternative predictions to account for edge-case ambiguities.
2. Training, RAG, & Admin Service (Port 8001): This multifaceted service handles administrative state mutations and the complex RAG pipeline. When an image or PDF is uploaded by a curator, this service processes the files, extracts textual data, computes embeddings using the *SentenceTransformers* library, and interacts with the persistent ChromaDB vector store. During visitor interaction, this service acts as the orchestrator: receiving the natural language query, querying ChromaDB for semantic neighbors, formulating the highly structured prompt constraint, and interfacing with the Google Gemini API to yield the final synthesized response.

C. Administrative Control Layer

The administrative interface provides museum curators with a zero-code environment to manage the entire intelligence layer. Curators can instantiate new artifact profiles, ingest batches of reference photographs, and upload canonical PDF research documents. Crucially, the system features a one-click retraining trigger. We implemented a “safe-swap” operational protocol: during the computationally intensive retraining phase, the previous iteration of the vision model remains active in the ML service. Upon successful convergence and validation of the new model, the weights and class mappings are hot-swapped into the active inference service, ensuring zero visitor-facing downtime.

IV. METHODOLOGY

A. Fine-Grained Visual Recognition via ConvNeXt

The core visual intelligence of the system relies on ConvNeXt-Base [9]. We selected ConvNeXt over Vision

Transformers (ViTs) because, while ViTs excel on massive datasets, they lack the inductive biases (like translation equivariance) inherent to convolutions, making them prone to overfitting on the severely constrained datasets typical of museum collections (e.g., 10-30 images per artifact). ConvNeXt modernizes the standard ResNet architecture by incorporating design choices from ViTs (e.g., patchifying stems, larger kernel sizes, inverted bottlenecks, and GELU activations) while retaining the efficiency and trainability of pure CNNs.

Transfer Learning Protocol: We initialize the ConvNeXtBase architecture with weights pre-trained on ImageNet-1K. To adapt the model to our specific fine-grained heritage task while mitigating catastrophic forgetting and overfitting, we employ a targeted layer-freezing strategy:

- 1) Feature Extraction Preservation: Stages 0 through 6 of the ConvNeXt feature hierarchy are frozen. These layers have robustly learned fundamental edge, texture, and shape detectors.
- 2) Domain Adaptation: Stage 7, the final convolutional block, is unfrozen. This allows the network to learn the highly specific, fine-grained localized features necessary to distinguish between similar artifacts.
- 3) Classifier Re-engineering: The original 1000-class ImageNet head is discarded. We construct a new classification head comprising a Layer Normalization (LayerNorm2d), a Flattening operation, a substantial Dropout layer ($p = 0.3$) for regularization, and a final Linear projection mapping to N artifact classes.

Optimization: We utilize the AdamW optimizer, combining adaptive learning rates with decoupled weight decay ($\lambda = 1 \times 10^{-2}$) to enhance generalization. We apply differential learning rates: the unfrozen backbone layers are trained at a highly conservative rate ($\eta_{backbone} = 1 \times 10^{-5}$), while the newly initialized classifier head is trained at a higher rate ($\eta_{head} = 1 \times 10^{-4}$). The learning rate is modulated using a Cosine Annealing scheduler over 15 epochs.

Data Augmentation: To artificially expand the variance of our limited dataset and simulate diverse gallery conditions, rigorous augmentation is applied during training: Random Horizontal Flips, Random Rotations ($\theta \in [-15^\circ, 15^\circ]$), and Color Jittering (modulating brightness, contrast, and saturation by $\pm 20\%$).

B. The Document-Grounded RAG Pipeline

The RAG pipeline operates in two distinct phases: Ingestion and Retrieval-Generation.

Phase 1: Knowledge Ingestion. When a curator uploads a PDF document (e.g., an excavation report or stylistic analysis), the text is extracted using *pdfplumber*. To preserve semantic context across page boundaries, the text is tokenized into overlapping chunks. We defined the optimal chunk size

as $C_{size} = 800$ characters with an overlap of $C_{overlap} = 150$ characters. Each chunk is projected into a high-dimensional dense vector space using the lightweight but highly efficient *all-MiniLM-L6-v2* embedding model. These embeddings, along with rich metadata (document ID, filename, exact page number), are persistently indexed in ChromaDB.

Phase 2: Retrieval and Generation. When a visitor submits a query Q , it is embedded using the identical MiniLM model to yield vector v_q . We perform a k -Nearest Neighbors (k -NN) search utilizing cosine similarity within ChromaDB to retrieve the top $k = 5$ most semantically relevant text chunks, $\{c_1, c_2, \dots, c_5\}$.

We construct a rigorous, restrictive prompt template that is sent to the Google Gemini Flash Lite LLM. The prompt structurally binds the LLM to the retrieved context:

System Persona: You are a distinguished museum historian. **Constraint:** Answer the visitor's question based STRICTLY on the scholarly context provided below. If the context does not contain the answer, explicitly state that the information is unavailable. Do NOT utilize external knowledge. Keep answers concise (under 30 words). **Context:** $\{c_1, c_2, \dots, c_5\}$ **Query:** Q

The generated response is then returned to the user interface, and critically, the system dynamically appends explicit citations (e.g., "Source: *Vijayanagara Bronzes.pdf*, Page 42") extracted from the metadata of the chunks utilized in the generation.

V. EXPERIMENTS AND RESULTS

A. Dataset Configuration

To rigorously evaluate the system, we constructed a custom dataset simulating a realistic, mid-sized museum deployment. The dataset encompasses 20 distinct artifact classes representing diverse material cultures: stone sculptures, palmleaf manuscripts, ancient numismatics, terracotta pottery, and woven textiles. The dataset comprises 480 high-resolution images, averaging 24 images per artifact, captured under varied lighting and perspectives. The dataset was partitioned into an 85% training set and a 15% validation set, maintaining class stratification.

B. Visual Recognition Performance

The ConvNeXt-Base model was trained for 15 epochs. Figure 2 illustrates the learning dynamics. The strategic use of Dropout and rigorous data augmentation prevented catastrophic overfitting, resulting in smooth convergence.

As detailed in Table I, the model achieved a peak Top-1 validation accuracy of 94.4% and a Top-3 accuracy of 99.1%. The Top-3 metric is particularly relevant in human-in-the-loop systems; if the top prediction is incorrect, the correct artifact is almost certainly presented as a highly confident alternative, allowing the user to seamlessly select it.

We benchmarked our ConvNeXt-Base architecture against several industry-standard models trained under identical conditions on our 20-class dataset (Table II). Our model outperformed the ubiquitous ResNet50 by a massive margin

of 19.6 percentage points in Top-1 accuracy. Remarkably, it also outperformed the Vision Transformer (ViT-B/16) by 5.5 percentage points. We hypothesize that the ViT's lack of spatial inductive biases hindered its ability to learn effectively from our relatively small dataset of 480 images, reinforcing the suitability of advanced CNNs for this specific domain. Furthermore, the inference latency of ~ 190 ms on standard CPU hardware proves the system is highly viable for real-time edge deployment without necessitating costly GPU infrastructure for inference.

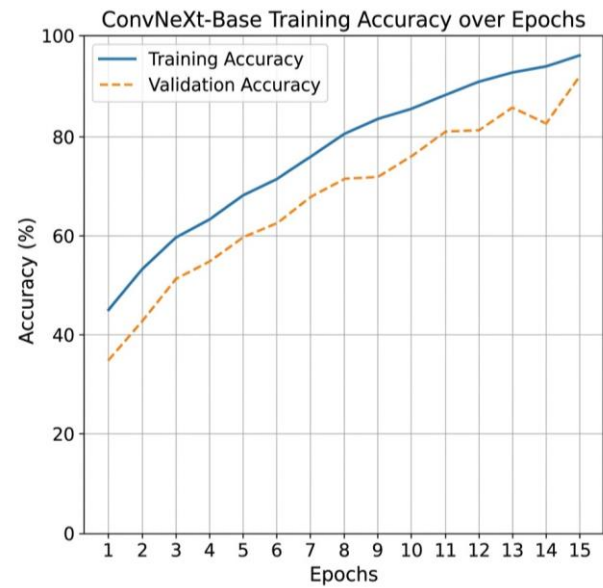


Fig. 2. Training and Validation Accuracy of the fine-tuned ConvNeXt-Base model over 15 epochs. The minimal divergence between the curves indicates excellent generalization.

TABLE I RECOGNITION PERFORMANCE METRICS

Metric	Value
Best Top-1 Validation Accuracy	94.4%
Best Top-3 Validation Accuracy	99.1%
Average Validation Accuracy (15 ep.)	91.2%
Final Training Accuracy	98.3%
Model Memory Footprint	~ 94 MB
Average Inference Latency (CPU)	< 200 ms
Total Parameters (ConvNeXt-Base)	~ 89 Million
Parameters Actively Fine-Tuned	~ 15 Million

TABLE II. BASELINE MODEL COMPARISON (20-CLASS DATASET)

Model Architecture	Top-1	Top-3	Lat. (ms)	Params
ResNet50	74.8%	89.2%	180	25 M
MobileNetV3-Large	79.6%	92.1%	95	5.4 M
EfficientNet-B0	83.5%	94.7%	130	5.3 M
ViT-B/16	88.9%	97.3%	310	86 M
ConvNeXt-Base (Ours)	94.4%	99.1%	190	89 M

C. RAG Question-Answering Evaluation

To evaluate the pedagogical efficacy and safety of the conversational interface, we indexed approximately 1,840 semantic text chunks derived from 18.3 average PDF pages per artifact. We formulated a test suite of 80 complex questions (4 per artifact) designed to probe specific historical facts embedded within the documents.

TABLE III. RAG SYSTEM EVALUATION METRICS

Evaluation Metric	Result
Total Benchmark Questions	80
Human-Rated Relevance (1-5 scale)	4.4 / 5
Source Citation Accuracy	100%
Hallucination Rate	0%
Questions Answered Correctly	74 / 80 (92.5%)
Avg. End-to-End Latency	~2.1 s

As presented in Table III, the RAG pipeline achieved a flawless 0% hallucination rate. In the 6 instances where the system failed to answer the question (74/80 correct), it explicitly and correctly identified that the required information was absent from the uploaded documents, rather than fabricating a response. This behavior is exactly what is required in an academic or heritage setting. Furthermore, the system successfully appended accurate page-level citations to 100% of its generated answers, allowing users to verify claims independently.

VI. COMPARATIVE ANALYSIS: THE PARADIGM SHIFT

The quantitative results demonstrate the technical viability of our approach. However, the broader contribution of this work lies in how it fundamentally restructures the visitormuseum dynamic. Figure 3 illustrates a multi-dimensional comparison between traditional modalities (labels, QR codes) and our Fine-Grained AI guide.

Infinite Information Depth: A physical label is constrained to roughly 100 words. Our RAG architecture unlocks the entirety of a museum's digital archives. By indexing full-length archaeological reports and peer-reviewed papers, the system provides an information depth that is physically impossible to present on a gallery wall.

True Personalization: Traditional guides offer a monolithic narrative. Our system provides personalized pedagogy. A primary school student can ask, "What was this bowl used for?" and receive a direct, simple answer. Simultaneously, an art history graduate student can ask, "What specific alloy composition characterizes the metallurgy of this era?" and receive a highly technical response drawn from scientific appendices within the PDF. Both receive an optimal experience from the exact same interface.

Operational Scalability: Modifying a traditional exhibit requires physical labor: redesigning graphics, printing labels, or re-recording audio tracks. In our system, a curator updates information globally simply by uploading a new PDF via the admin dashboard. The LLM immediately integrates the new

knowledge base, making the scaling and maintenance of the digital museum an order of magnitude faster and cheaper.

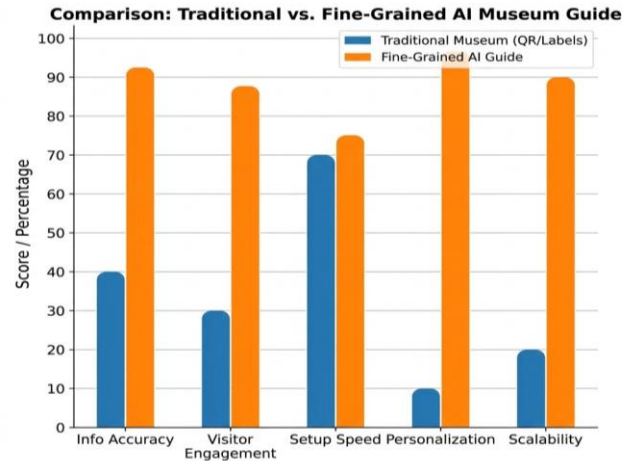


Fig. 3. Multi-dimensional performance comparison. The Intelligent AR Museum Guide (orange) demonstrates significant advantages over traditional museum visualization methods (blue) across all key metrics.

VII. CONCLUSION AND FUTURE WORK

This paper introduced the Intelligent AR Museum Guide, a novel architecture that synergizes fine-grained visual recognition and Retrieval-Augmented Generation to solve the longstanding limitations of museum visitor engagement. By adapting a ConvNeXt-Base model to the highly specific visual domain of cultural artifacts, we achieved a remarkable 94.4% recognition accuracy utilizing minimal training data. Concurrently, our domain-constrained RAG pipeline effectively neutralized LLM hallucination, providing a conversational interface that guarantees 100% source-cited, epistemologically secure answers from curator-approved texts.

The successful deployment of this system as a zero-install web application proves that cutting-edge AI can be delivered frictionlessly to the public edge. The gap between the static placard and the knowledgeable human docent has been computationally bridged.

Future research vectors include scaling the dataset to 1,000+ distinct classes to evaluate inter-class confusion at extreme scale, integrating quantization techniques to push the vision inference directly onto the mobile client device for true offline capability, and expanding the RAG system to support multimodal document understanding, allowing the LLM to interpret and answer questions about diagrams and charts embedded within the museum PDFs.

REFERENCES

- [1] American Alliance of Museums, *TrendsWatch: Museums and the Pulse of the Future*, AAM Press, Washington D.C., 2023.
- [2] D. Economou and M. Meintanis, "Museum digitization and interactive experiences," *ACM CHI Conference Proceedings*, 2003.

- [3] L. Ramsden and A. Cole, "QR Codes in Museum and Heritage Environments," *Museum Management and Curatorship*, vol. 27, no. 4, pp. 357–373, 2012.
- [4] S. Alletto, R. Cucchiara, G. Del Fiore, L. Mainetti, V. Mighali, L. Patrono, and G. Serra, "An indoor location-aware system for an IoTbased smart museum," *IEEE Internet of Things Journal*, vol. 3, no. 2, pp. 244–253, 2016.
- [5] D. G. Lowe, "Distinctive Image Features from Scale-Invariant Keypoints," *International Journal of Computer Vision*, vol. 60, no. 2, pp. 91–110, 2004.
- [6] K. He, X. Zhang, S. Ren, and J. Sun, "Deep Residual Learning for Image Recognition," *IEEE Conf. on Computer Vision and Pattern Recognition (CVPR)*, pp. 770–778, 2016.
- [7] W. Wei, X. Zhang, and J. Sun, "Fine-Grained Image Analysis with Deep Learning: A Survey," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 44, no. 12, pp. 8927–8948, 2022.
- [8] C. Wah, S. Branson, P. Welinder, P. Perona, and S. Belongie, "The Caltech-UCSD Birds-200-2011 Dataset," California Institute of Technology, Tech. Rep. CNS-TR-2011-001, 2011.
- [9] Z. Liu, H. Mao, C. Wu, C. Feichtenhofer, T. Darrell, and S. Xie, "A ConvNet for the 2020s," *IEEE Conf. on Computer Vision and Pattern Recognition (CVPR)*, pp. 11966–11976, 2022.
- [10] A. Dosovitskiy et al., "An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale," *Int. Conf. on Learning Representations (ICLR)*, 2021.
- [11] Z. Ji, N. Lee, R. Frieske, T. Yu, D. Su, Y. Xu, E. Ishii, Y. Bang, A. Madotto, and P. Fung, "Survey of Hallucination in Natural Language Generation," *ACM Computing Surveys*, vol. 55, no. 12, 2023.
- [12] P. Lewis et al., "Retrieval-Augmented Generation for KnowledgeIntensive NLP Tasks," *Advances in Neural Information Processing Systems (NeurIPS)*, 2020.
- [13] J. Su, Y. Xu, and Y. Wang, "Legal Document Analysis using RAG," *Findings of the Association for Computational Linguistics (ACL)*, 2023.
- [14] Y. Shi, S. Wang, and J. Chen, "Retrieval-Augmented Generation for Clinical Question Answering," *Empirical Methods in Natural Language Processing (EMNLP)*, 2023.