

MedVisionEx: A Data-Information-Knowledge-Wisdom Hierarchy Inspired, Retrieval-Augmented and Knowledge-Graph-Grounded Pipeline for Explainable Medical Visual Question Answering

Nadeesha Perera¹  Alli Raittinen² 

Jyrki Nummenmaa³ 

Konstantinos Stefanidis⁴ 

Tampere University, Tampere, Finland

nadeesha.perera@tuni.fi alli.raittinen@tuni.fi jyrki.nummenmaa@tuni.fi
konstantinos.stefanidis@tuni.fi

Abstract. Medical visual question answering (VQA) systems are rarely adopted in practice because they answer as opaque black boxes: a clinician sees a verdict but neither the visual evidence behind it nor the medical knowledge that supports it. In radiology diagnostics, accuracy alone is insufficient; clinicians need justification. We demonstrate *MedVisionEx*, an on-premise pipeline whose design follows the Data-Information-Knowledge-Wisdom (DIKW) hierarchy: raw *data* (image and question) is turned into *information* by retrieving similar archived cases, enriched into *knowledge* by linking to a biomedical knowledge graph, and finally elevated to *wisdom*—a grounded, traceable answer. Concretely, a query image is encoded by a domain-adapted CLIP encoder, fine-tuned with a hard-negative contrastive objective over PCA-whitened embeddings, and used to retrieve visually similar cases from a FAISS index pooled over several public radiology collections. The retrieved images and captions condition a local finetuned vision-language model based on MedGemma that writes a concise, evidence-grounded description; this description and the question are linked to UMLS concepts (subsuming SNOMED CT and RadLex), and the resulting structured knowledge augments both the description and the question-answering prompt. A generative answer model produces the final answer. Retrieval supplies reliability from real data; the knowledge graph supplies explainability from structured relations; together they yield grounded outputs with reduced hallucinations, and traceable reasoning. Every component runs locally, and the interface exposes the retrieved cases, the generated description, and the knowledge-graph support so users can audit how each answer was reached. We describe MedVisionEx as a novel DIKW pyramid-inspired pipeline architecture, that combines known retrieval augmentation with FAISS databases and knowledge graphs, the data-management design behind the index, and how the pipeline establishes explainability and reliability necessary in medical AI.

Keywords: Medical VQA DIKW hierarchy Retrieval-augmented generation Knowledge graphs
Vision-language models Explainable AI.

1 Introduction

Medical AI research in radiology diagnostics is expanding rapidly, yet two obstacles keep such systems out of the clinical setting. First, end-to-end VQA models are *black boxes*: they produce an answer without analyzing the image regions to comparable cases, or grounding on domain facts that justify it. In a setting where every statement must be explainable, accuracy alone is not enough—clinicians need justification, because a confident but unexplained answer is, in practice, unusable for a decision that carries clinical risk. Second, deployment constraints in hospitals frequently rule out sending images to external APIs, so a practical system must run *on-premise* on commodity accelerators. These two pressures—the need for justification and the need for local execution—jointly rule out the dominant pattern of querying a large, remote, monolithic model, and motivate a transparent composition of smaller local components whose intermediate outputs can be inspected.

We demonstrate *MedVisionEx*, a system designed around the principle that an answer is only as useful as the evidence a user can inspect behind it. Rather than mapping image and question directly to an answer, the pipeline first *retrieves* similar archived cases, then *writes* an evidence-grounded description of the query image conditioned on those cases, then *grounds* both the description and the question in a biomedical knowledge graph, and only then provide *answers*. Each intermediate step—the retrieved cases with similarity scores, the generated description, and the linked concepts and relations—is preserved and shown to the user, producing a traceable path from image to answer.

This staged design is not incidental: it instantiates the Data-Information-Knowledge-Wisdom (DIKW) hierarchy [1], a classic model of how raw data becomes actionable insight. Each pipeline stage occupies a level of the pyramid (Sect. 2), and two mechanisms drive the ascent; (1) retrieval, which grounds the system in real *Information* is the source of *reliability*, and (2) the knowledge graph, which provides structured domain *Knowledge* is the source of *explainability*. Their combination yields grounded outputs, fewer hallucinations, and reasoning a clinician can follow.

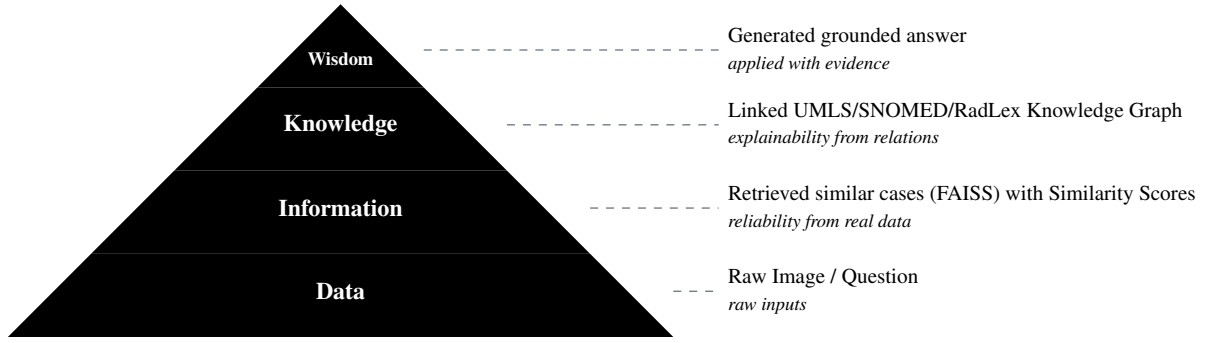


Figure 1: The DIKW ascent realized by MedVisionEx. Each rung is a concrete, inspectable pipeline stage. Retrieval lifts data to information and is the source of reliability; the knowledge graph lifts information to knowledge and is the source of explainability; the answer model applies that knowledge to produce wisdom.

From a data-management perspective the core asset is the similarity index pooled across heterogeneous public radiology datasets, built on a representation learned specifically to make clinically similar image neighbours. The database pipeline manages and mines heterogeneous data for building the data services for our XAI (Explainable AI) application. Our contributions are: (i) a DIKW-inspired, retrieval-augmented, knowledge-graph-grounded VQA architecture that is explainable by construction; (ii) a pooled multi-source similarity index with a contrastively fine-tuned, whitened embedding space; (iii) a knowledge-graph integration over UMLS, SNOMED CT, and RadLex with a dependency-light linking back end that runs without heavy native libraries, easing on-premise deployment; and (iv) a decoupled, auditable service—an inference API plus a thin client—whose interactive front end exposes the full evidence chain.

The paper details the DIKW design rationale (Sect. 2), the architecture (Sect. 3), knowledge-graph integration (Sect. 4), implementation and data management (Sect. 5), the demonstration of examples (Sect. 6), and an evaluation protocol (Sect. 7).

2 Design Rationale: From Data to Wisdom

The DIKW hierarchy describes how raw *data* gains context to become *information*, how information gains meaning and relations to become *knowledge*, and how knowledge is applied with judgment to yield *wisdom*. MedVisionEx was designed so that each layer corresponds to a concrete, inspectable stage, and so that the two persistent weaknesses of medical generative AI—unreliability and opacity—are addressed at the layers where they arise. Figure 1 summarizes the DIKW mapping.

The value of the framing is operational, not merely descriptive. Because the *information* layer is designed as the retrieved cases, every answer can be traced to specific archived evidence; because the *knowledge* layer is developed as the linked concepts and relations, the domain facts behind an answer are explicit rather than buried in model weights. The two layers attack hallucination from complementary directions: retrieval constrains the description to be consistent with real cases, and the knowledge graph constrains the answer to be consistent with established relations between findings, anatomy, and diseases.

2.1 Reliability and Explainability as Complementary Mechanisms

It is worth separating the two properties because they are produced by different mechanisms and fail in different ways. *Reliability* is the degree to which an output is anchored in real, verifiable evidence rather than fabricated; in MedVisionEx it comes from retrieval, because the description is conditioned on archived cases whose captions are ground-truth clinical text. When retrieval returns strong neighbours the description has external support; when it returns nothing above threshold, the system makes the absence of evidence visible by falling back to a zero-shot description rather than silently guessing. *Explainability* is the degree to which a user can follow *why* an output is valid; in MedVisionEx it comes from the knowledge graph, because the concepts and relations bridging the question and the description are surfaced explicitly. A purely retrieval-based system can be reliable yet opaque about the clinical reasoning, and a purely knowledge-based system can be explainable yet contain no context from the specific image at hand. Combining them is what produces outputs that are simultaneously grounded and interpretable. This separation also guides evaluation (Sect. 7): retrieval coverage probes reliability, while the knowledge-graph ablation probes explainability and its downstream effect on answer quality.

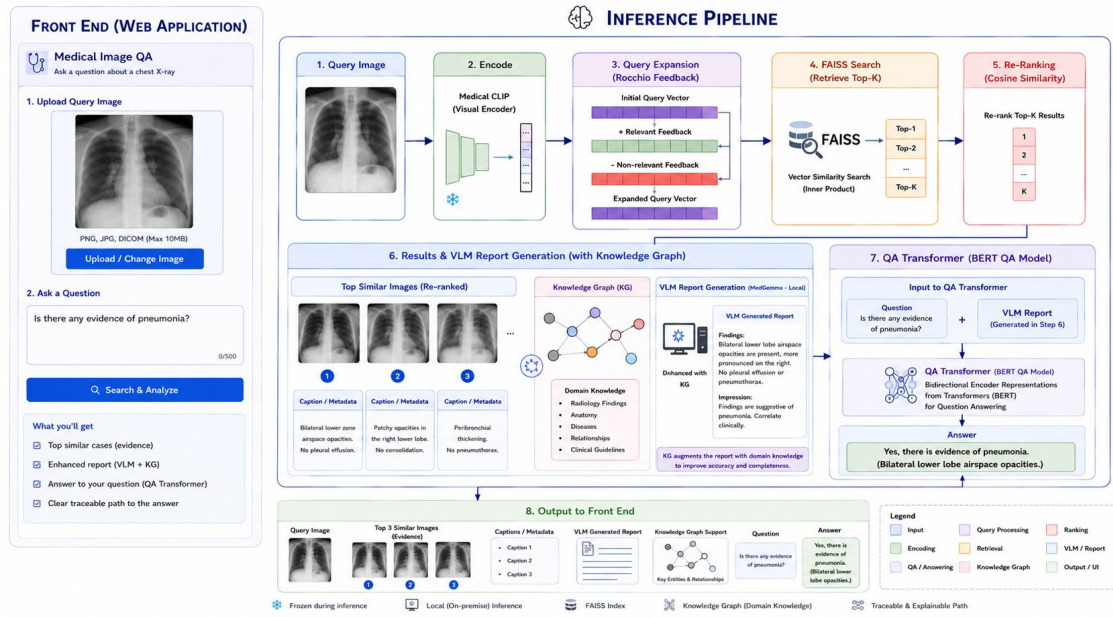


Figure 2: The MedVisionEx architecture. A query image is encoded and used to retrieve similar cases from a pooled FAISS index; the retrieved images and captions condition a local vision–language model that writes an evidence-grounded description, which a knowledge graph augments; a generative answer model then answers the question from the description, the knowledge context, and the question. The front end (left) exposes every intermediate artifact, yielding a traceable path from image to answer.

3 System Overview

Figure 2 shows the complete pipeline and the web front end. Given a query image and a question, processing proceeds in eight stages.

(1) Query image and (2) encoding [Data]. The uploaded image is embedded by a medical CLIP visual encoder whose weights are frozen at inference. The encoder is fine-tuned offline (Sect. 5) so that images sharing clinical content lie close in the embedding space.

(3) Query refinement. The query embedding may be refined by pseudo-relevance feedback, nudging it toward confidently retrieved neighbours and away from dissimilar ones before the main search. This step is optional and disabled by default in the demonstration.

(4) Similarity search [Information]. The embedding is matched against a FAISS index [5] by inner product. Rather than a fixed k , the system retrieves *all* gallery items whose similarity exceeds a threshold, so the amount of evidence scales with how well the case is covered by the archive; when nothing clears the threshold the pipeline degrades gracefully to a zero-shot description.

(5) Ranking. Retrieved candidates are ordered by cosine similarity and capped for presentation; their scores travel with them so the interface can display how strongly each case supports the answer.

(6) Evidence-grounded description with knowledge graph [Knowledge]. The query image, together with the retrieved images and their captions, is passed to a local finetuned MedGemma [12] model that writes a concise single-paragraph description stating the modality, the imaged anatomy, and any abnormalities. A biomedical entity linker extracts the findings present in the retrieved captions and injects them as structured hints into the prompt (Sect. 4).

(7) Answering [Wisdom]. The generated description and the question are linked to knowledge-graph concepts; the concepts, their definitions, and relations connecting question terms to description terms are serialized into a compact knowledge block. A generative answer model consumes the description, this knowledge block, and the question to produce the final answer.

(8) Output. The front end returns the answer alongside the query image, the top similar cases, their captions, the generated description, and the knowledge-graph support, making the reasoning path auditable.

4 Where to Integrate the Knowledge Graph

A distinguishing design choice is *where* domain knowledge enters. MedVisionEx links text to the Unified Medical Language System (UMLS) [8], which incorporates both SNOMED CT [11] and RadLex [10], so a single linker [9] provides concepts, semantic types, definitions, synonyms, and—through a local relation table—ontology relations across all three vocabularies. The knowledge graph encodes the clinically

meaningful links identified as the origin of explainability: diseases to symptoms, and anatomy to findings. Knowledge is extracted once from the question and the generated description and reused at two points.

Caption side. Findings detected in the retrieved captions (diseases, signs, abnormalities) are passed to the description model as context, encouraging it to confirm or refute clinically plausible findings rather than overlook them, while the instruction still directs the model to rely on what is visible.

Answer side. This is where grounding most directly helps the answer. The question and the generated description are linked; the system produces the standard names and semantic types of the mentioned concepts, short definitions of key findings, the concepts shared between question and description, and any one-hop ontology relations connections. This block is attached to the answer model’s prompt, supplying definitions and relations that may be present in the image but absent from the description. Because the answer model is instruction-tuned, this requires no retraining; the same block is fine-tuned to rely on the knowledge.

Modality awareness. A recurring difficulty in pooled medical retrieval is that the imaging *modality* (radiograph, CT, MRI, ultrasound) is itself a major axis of similarity and confusion. The knowledge graph is therefore used as a guide to identify modality and disease evidence in particular, so that the description and answer remain connected to the correct modality rather than drifting towards a visually similar but modality-mismatched neighbour.

5 Implementation and Data Management

Pooled similarity index. The retrieval gallery unifies several public radiology image–caption collections [14, 15, 16] into one corpus. Each record is normalized to an (image, caption) pair and sources that expose only a training split contribute a deterministic held-out slice so the index can be evaluated without leakage. The pooled embeddings are indexed with FAISS under inner product. Because indexing and search are decoupled from training, the gallery can be extended incrementally: validation embeddings can be appended to the trained index without rebuilding it.

Contrastive fine-tuning. The CLIP [6] encoder is adapted with a symmetric InfoNCE [13] (Information Noise Contrastive Estimation) objective strengthened by a hard-negative loss term. The hardest negatives are mined within each batch and from recent embeddings, expanding the effective negative pool without extra encoder passes; mining is enabled only after a warm-up so that the embedding space stabilizes first. This sharpens the contrast between visually similar but clinically different studies, which is exactly the confusion that reduces evidence quality and, with it, reliability.

Embedding post-processing. Before indexing, image embeddings are decorrelated and scaled by a PCA whitening transform; the identical transform is applied to query embeddings at search time. Whitening removes dominant shared directions that otherwise increase cosine similarity, improving the separation of true neighbours from generic matches.

Local generation. Both the description model and the answer model run on-premise, keeping data within the institution. The description model receives the query image and up to a bounded number of retrieved images so that prompt size and accelerator memory stay controlled, while all retrieved captions are available as text. The answer model is parameter-efficiently fine-tuned to map a description and a question to a short answer; conditioning the answer on the generated description, rather than on the raw image alone, is what lets retrieved evidence inform the response.

Knowledge linking. Entity linking supports two interchangeable back ends behind one interface. A neural back end uses a biomedical NLP pipeline [9] with abbreviation resolution and a UMLS linker. A second, dependency-light back end performs greedy longest-match linking against a compact concept dictionary compiled offline from the UMLS subset, keyed by normalized term to concept identifier with semantic-type filtering; it uses only the standard library and so loads without the heavy native dependencies the neural linker requires—a practical advantage for hardened on-premise hosts. Even if the dictionary is unavailable, knowledge augmentation can be skipped such that the pipeline still answers. Linked concepts are cached by text so that repeated descriptions across multiple questions about the same image are not re-processed, and relations are read from a local ontology table so no external service is contacted.

Reproducibility. The pipeline is configuration-driven: dataset specifications, the similarity threshold, the number of shown evidence images, generation lengths, and knowledge-graph depth are all declared in configuration, and each stage can be run independently (train → index → append → evaluate).

Deployment. The system is delivered as a decoupled service: a stateless inference API loads the encoder, index, description model, knowledge linker, and answer model once on the accelerator host and exposes a single endpoint that accepts an image (including DICOM) and a question and returns the answer together with every intermediate step—the retrieved cases and scores, the generated description, and the knowledge-graph support—serialized as one auditable response. A separate thin web client renders this response and runs on any machine, so the graphical interface never holds model weights or patient data on disk. Inference is serialized on the single accelerator, and throughput scales by replicating the API rather than by sharing

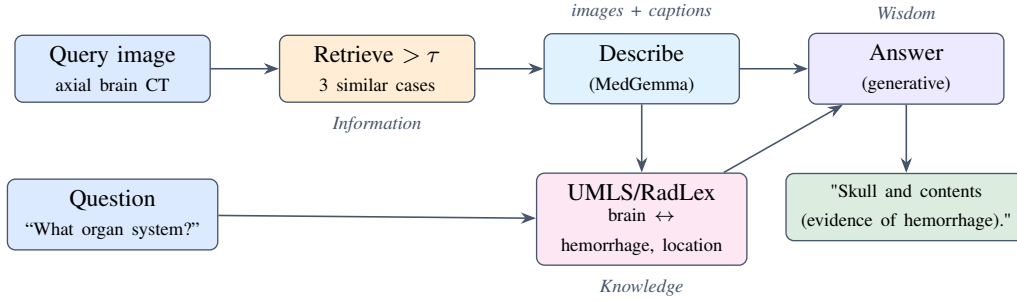


Figure 3: Worked example annotated with DIKW layers using a brain CT query from 4 Panels D, E, F: an axial brain image and anatomical question (data) flow through retrieval of similar intracranial cases (information), a case-conditioned description of findings including hemorrhage, knowledge-graph linking of brain anatomy, imaging modality, and pathology (knowledge), and generative organ-system identification with clinical grounding (wisdom). Each box corresponds to an artifact the interface exposes.

Table 1: End-to-end evaluation on four medical VQA benchmarks: VQA-RAD [17], SLAKE [18], Image-CLEF 2019 Med-VQA [19], and RadImageNet-VQA [20]. BERTScore-F1 [2] measures semantic similarity to ground-truth answers; ROUGE [3] captures reference overlap; Accuracy measure answer correctness and clinical relevance.

Dataset	BERTScore-F1	ROUGE-L	Accuracy
VQA-RAD	0.962	0.531	0.531
SLAKE	0.939	0.649	0.641
Med-VQA 2019	0.975	0.687	0.687
RadImageNet	0.964	0.63	0.63

one model across threads.

6 Demonstration

Figure 3 traces a single query end to end, annotated with the DIKW layer each step occupies, taken from the example in D,E,F panels of 4. The user uploads an axial brain CT scan and asks, “What organ system is shown?” (data). Retrieval returns archived cases whose captions describe non-contrast brain CT imaging and intracranial findings (information); the description model, conditioned on these cases and the image, writes a paragraph noting the axial brain CT appearance with evidence of hemorrhage in the right basal ganglia region and surrounding edema. The knowledge graph links *brain* (from the visual findings) and *hemorrhage* (from the description) and supplies relations such as anatomical location (*head-brain*), imaging modality (*CT scan of brain*), and associated pathophysiology (*parenchyma, edema*) (knowledge). The answer model, given the description, this knowledge, and the question, responds “skull and contents” identifying the organ system while grounding the finding in the retrieved similar cases and linked clinical concepts (wisdom)—and the interface shows the three supporting cases and the knowledge graph with anatomical and pathological relationships beside the answer.

7 Evaluation Protocol

We evaluate the full pipeline on a held-out set of image–question–answer triples. For each pair the system executes the entire chain—retrieve, describe, ground, answer—and the predicted answer is compared with the reference answer. Because answers are free text, we report BERTScore [2], ROUGE [3] and additionally accuracy as the metrics standard for generative outputs; descriptions are cached per image since a set typically poses several questions about the same study.

The knowledge graph can be switched off to isolate its contribution to answer quality, directly testing the deck’s claim that knowledge grounding improves medical VQA; the similarity threshold can be varied to trade evidence breadth against precision. Beyond answer similarity we report *retrieval coverage*—the fraction of queries that find above-threshold evidence—because it determines how often an answer is grounded in real cases (the information layer) versus produced zero-shot, and is thus a direct measure of reliability.

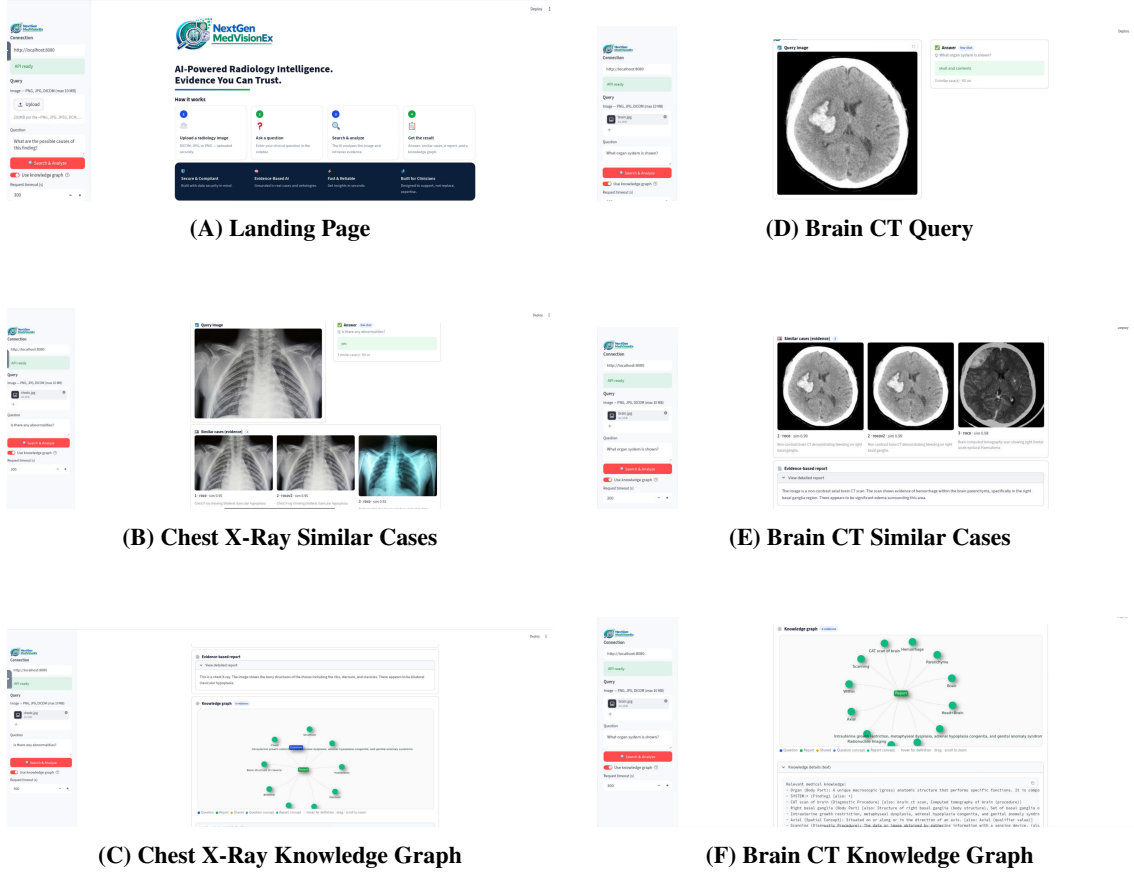


Figure 4: **Demonstrations from the MedVisionEx Application.** Panel (A) shows the landing page with the application interface. Panels (B) and (C) display query results for a chest X-ray image, including similar cases retrieved from the database and their associated knowledge graph with medical concepts and relationships. Panels (D), (E), and (F) demonstrate the same workflow applied to a brain CT scan, showcasing the system’s ability to retrieve semantically similar cases and extract relevant clinical knowledge for diagnosis support.

8 Related Work

MedVisionEx draws on three threads. Retrieval-augmented generation grounds generative models in retrieved evidence [4]; we apply it to images, retrieving comparable cases before description, which supplies the information layer. Vision–language models for radiology generate reports from images [7]; we condition such a model on retrieved evidence and run it locally. Knowledge graphs built on UMLS, SNOMED CT, and RadLex have grounded clinical text understanding; we use them to augment both description and answering, supplying the knowledge layer. The DIKW hierarchy [1] has long informed information-systems design; our contribution is to operationalize it as a concrete, inspectable medical-VQA pipeline in which every layer is a first-class output and the similarity gallery is treated as a managed, extensible data service. In contrast to monolithic medical VQA models, the composition is explainable by construction.

9 Limitations and Future Vision

The present system has clear limitations that also chart its roadmap. The most pressing is *modality*: when a single index pools radiographs, CT, MRI, and ultrasound, the encoder can retrieve a visually similar neighbour of the wrong modality, and a description grounded in such a neighbour inherits that error. We currently mitigate this with the contrastive objective and by steering the knowledge graph toward modality evidence, but making the retrieval representation explicitly modality-aware—and verifying modality consistency between query and retrieved cases before they condition the description—is a natural next step. A second limitation is that knowledge enters at inference only; folding the serialized knowledge block into the answer model’s fine-tuning data would teach it to rely on the graph rather than merely tolerate it. A third is coverage: the relation component depends on a local ontology subset, so relations are only as rich as the table provided, and out-of-distribution query images (from a different acquisition protocol than the gallery)

reduce retrieval coverage and push the system toward zero-shot descriptions.

The broader vision is the trajectory from black-box prediction toward trustworthy clinical AI. Reliability and explainability are not add-ons but the organizing goals: an output a clinician can trace to real cases and established relations is one they can act on or challenge. By treating the evidence gallery as a managed, extensible data service and the knowledge graph as a first-class reasoning substrate, MedVisionEx is positioned to grow toward clinical alignment—incorporating more modalities, richer relations, and tighter coupling between retrieval, knowledge, and generation—without surrendering the auditability that motivates the design.

10 Conclusion

We presented MedVisionEx, an on-premise medical VQA pipeline whose architecture follows the DIKW ascent: retrieval turns data into information and supplies reliability, a knowledge graph turns information into knowledge and supplies explainability, and a local generative model turns knowledge into a grounded, traceable answer. By pooling heterogeneous radiology collections into one contrastively learned, whitened similarity index and grounding generation in both retrieved evidence and biomedical ontologies, the system reduces hallucination and makes its reasoning auditable end to end. The demonstration lets attendees manipulate the evidence pathway directly and observe its effect on answers, illustrating a data-centric route from black-box prediction toward trustworthy, clinically aligned decision support.

10.0.1 Disclosure of Interests The authors have no competing interests to declare that are relevant to the content of this article.

References

- [1] Rowley, J.: The wisdom hierarchy: Representations of the DIKW hierarchy. *J. Inf. Sci.* **33**(2), 163–180 (2007)
- [2] Zhang, T., Kishore, V., Wu, F., Weinberger, K.Q., Artzi, Y.: BERTScore: Evaluating text generation with BERT. In: *International Conference on Learning Representations (ICLR)* (2020)
- [3] Lin, C.-Y.: ROUGE: A package for automatic evaluation of summaries. In: *Text Summarization Branches Out*, pp. 74–81. Association for Computational Linguistics (2004)
- [4] Lewis, P., Perez, E., Piktus, A., Petroni, F., Karpukhin, V., Goyal, N., Küttler, H., Lewis, M., Yih, W.-t., Rocktäschel, T., Riedel, S., Kiela, D.: Retrieval-augmented generation for knowledge-intensive NLP tasks. In: *Advances in Neural Information Processing Systems (NeurIPS)*, vol. 33, pp. 9459–9474 (2020)
- [5] Johnson, J., Douze, M., Jégou, H.: Billion-scale similarity search with GPUs. *IEEE Trans. Big Data* **7**(3), 535–547 (2021)
- [6] Radford, A., Kim, J.W., Hallacy, C., Ramesh, A., Goh, G., Agarwal, S., Sastry, G., Askell, A., Mishkin, P., Clark, J., Krueger, G., Sutskever, I.: Learning transferable visual models from natural language supervision. In: *Proceedings of the 38th International Conference on Machine Learning (ICML)*, pp. 8748–8763 (2021)
- [7] Tu, T., et al.: Towards generalist biomedical AI. *NEJM AI* **1**(3) (2024)
- [8] Bodenreider, O.: The Unified Medical Language System (UMLS): Integrating biomedical terminology. *Nucleic Acids Res.* **32**(Database issue), D267–D270 (2004)
- [9] Neumann, M., King, D., Beltagy, I., Ammar, W.: ScispaCy: Fast and robust models for biomedical natural language processing. In: *Proceedings of the 18th BioNLP Workshop and Shared Task*, pp. 319–327 (2019)
- [10] Langlotz, C.P.: RadLex: A new method for indexing online educational materials. *RadioGraphics* **26**(6), 1595–1597 (2006)
- [11] Donnelly, K.: SNOMED CT: The advanced terminology and coding system for eHealth. *Stud. Health Technol. Inform.* **121**, 279–290 (2006)
- [12] Rajan, R., et al.: MedGemma: Open Models for Medical Text and Image Understanding. *arXiv preprint arXiv:2507.05201* (2025)

- [13] van den Oord, A., Li, Y., Vinyals, O.: Representation Learning with Contrastive Predictive Coding. In: *Advances in Neural Information Processing Systems (NeurIPS)*, vol. 31 (2018)
- [14] Pelka, O., Koitka, S., Rückert, J., Nensa, F., Friedrich, C.M.: Radiology Objects in COntext (ROCO): A multimodal image dataset. In: *Large-scale Annotation of Biomedical Data and Expert Label Synthesis (LABELS)*, *Lecture Notes in Computer Science*, vol. 11031, pp. 180–189. Springer (2018)
- [15] Pelka, O., Koitka, S., Nensa, F., Friedrich, C.M.: Overview of ImageCLEFmed 2020: Automatic generation of radiology reports, visual question answering, and caption prediction. In: *CLEF 2020 Working Notes, CEUR Workshop Proceedings* (2020)
- [16] Johnson, A.E.W., Pollard, T.J., Berkowitz, S.J., Greenbaum, N.R., Lungren, M.P., Deng, C.-Y., Mark, R.G., Horng, S.: MIMIC-CXR: A large publicly available database of labeled chest radiographs. *Sci. Data* **6**, 317 (2019)
- [17] Lau, J.J., Gayen, S., Abacha, A.B., Demner-Fushman, D.: A dataset of clinically generated visual questions and answers about radiology images. In: *Proceedings of the Scientific Document Understanding Workshop at AAAI* (2018)
- [18] Liu, B., Zhan, L.-M., Xu, L., Ma, L., Yang, Y., Wu, X.: SLAKE: A semantically-labeled knowledge-enhanced dataset for medical visual question answering. In: *IEEE International Symposium on Biomedical Imaging (ISBI)*, pp. 1650–1654 (2021)
- [19] Ben Abacha, A., Hasan, S.A., Datla, V., Liu, J., Demner-Fushman, D., Müller, H.: VQA-Med: Overview of the medical visual question answering task at ImageCLEF 2019. In: *CLEF 2019 Working Notes, CEUR Workshop Proceedings* (2019)
- [20] Butsanets, L., Corbière, C., Khlaout, J., Manceron, P., Dancette, C.: RadImageNet-VQA: A Large-Scale CT and MRI Dataset for Radiologic Visual Question Answering. *arXiv preprint arXiv:2512.17396* (2025)