



Deep Learning for Clinical Imaging Diagnosis at Scale

Luca Bianchi

Asst Professor

Abstract

Recent advances in deep learning have generated excitement within the domain of medical imaging diagnostics. Promising results achieved in specific tasks, often with limited training data, have led to calls for scaled-up applications leveraging large-scale medical imaging datasets. Such an approach could facilitate rapid and cost-effective development of diagnostic algorithms based on different imaging modalities. Nevertheless, a crucial step toward subsequent clinical uptake of these methods is the successful deployment of a large, multicenter test set for performance evaluation. However, despite being essential to the transfer of deep learning methods into clinical practice, scalable evaluation and deployment of these technologies have yet to receive significant attention. The attractions of large-scale testing and the potential effects of scale on accuracy are therefore often overlooked.

Analysis of deep learning methods — including convolutional neural networks, transformers, and other architectures — capable of utilizing large-scale clinical imaging datasets to achieve comparable or superior sensitivity, specificity, and area under the curve values to human clinicians is presented. Key strengths of the methodology include the emphasis on large-scale datasets, consideration of diagnostic performance across diverse clinical specialties, and application of external validation datasets to assess generalization performance across multiple sites. These qualities make the method particularly suitable for the evaluation of potentially biased or miscalibrated systems.

Keywords : Medical imaging; deep learning; imaging diagnoses; convolutional neural networks; transformer; multicenter clinical datasets ,Data Generation Metrics, Biases, and Capability, Generative AI Model Evaluation in Healthcare.

1. Introduction

Advances in medical imaging diagnostics offer the potential for earlier detection of many diseases and conditions, enable more effective therapeutic decision making, improve patient prognosis and quality of life, and reduce healthcare expenditures. However, for some critical diseases, notably cancer, mortality rates have remained static or even increased in the past few years, pointing towards deficiencies in clinical imaging practices. Misdiagnosis is widespread, whether due to detection, classification, or localization errors. Diagnostic accuracy also suffers from clinician fatigue, experience level, and availability. Such limitations in medical imaging diagnostics offer the potential for considerable clinical impact—if sufficiently large clinical datasets can be assembled—and will therefore serve as the focus of examination in this analysis.

Deep learning, developed for image-related tasks such as object detection, recognition, translation, and generation, offers a particularly promising means to address these challenges. With the availability of large-scale datasets, researchers have demonstrated that deep-learning-enabled systems can achieve state-of-the-art performance on many tasks. A system for CT image detection, classification, and localization of pulmonary nodules has been proposed, while earlier warning of malignancy



risk through classification of low-cancer-risk and high-cancer-risk patients has been realized with MRI images. Furthermore, under the supervision of a senior clinician, deep missense algorithms can aid junior radiologists in detecting subtle conditions.

1.1. Background and Significance

Medical imaging is capable of detecting an impressive range of conditions across multiple organ systems. It is among the most widely used forms of clinical diagnostics in the world—billions of imaging examinations are performed on millions of patients each year. Yet, despite limitless clinical interpretation and monitoring capabilities, the asymmetrical distribution of technical resources, skilled human expertise, and advanced diagnosis methods poses unique challenges to health systems and patients alike. Studies demonstrate that medical image diagnoses require significant time and effort to interpret and represent an additional workload, fatigue source, and mental stress burden for health professionals. Detection and classification errors also occur regularly, limiting the reliability of diagnoses.

Recent advancements in artificial intelligence (AI) and deep learning present scalable solutions capable of supporting clinical decision-making during medical image diagnostics. These technologies can automatically and accurately detect and classify diseases and pathologies across several imaging modalities, potentially complementing and enhancing expert human judgment. Nevertheless, their clinical adoption remains less extensive than anticipated. Deep learning models often rely on small-scale, single-center datasets for training, and the corresponding digital pipelines are not shared with the scientific community—an approach that contradicts the core principles of scientific research, which emphasize reproducibility and verification. Moreover, previously reported performance gains were achieved using single- or limited-center datasets, introducing reliability concerns. To address these issues, the curated and publicly available large-scale clinical datasets listed below are leveraged. Hypotheses are tested by developing and evaluating a diverse range of state-of-the-art and specialized deep learning models integrated into clinical workflows throughout the digital pipelines.



Fig 1: Deep Learning–Based Medical Imaging Diagnostics

1.2. Research design

The study objectives encompass a multifaceted evaluation of deep learning–based medical imaging diagnostics that are sufficiently scalable for application with the diverse patient populations typically encountered in clinical settings. Specifically, the technologies are scrutinized for their responsiveness toward diagnostic tasks within potentially underrepresented classes, such as rare diseases, and their calibration for estimated certainty. Together, these concerns aim for clinically relevant replicability in independent, external evaluations, consistent with the interpretation standards required for human–century partnerships. Two primary research questions guide the investigation: What large-scale datasets, representative of typical clinical practice, are publicly accessible for testing purposes? And how are large-scale clinical investigations convincingly executed?

Establishing the clinical relevance of a diagnostic model should rest on demonstrably reproducible testing with extensive patient cohorts sourced from multiple hospitals and medical centers. Accordingly, the promise of deep learning for medical imaging is explored using external validation within the massive Medical Information Mart for Intensive Care-III database. The data record 53 423 clinical encounters occurring in the Intensive Care Unit at the Beth Israel Deaconess Medical Center, Boston, Massachusetts, from 2001 to 2012. Eleven deep learning classification and regression models covering 56 distinct tasks are trained with data provenance and distribution intentionally concealed and subsequently evaluated by an independent third party. Scaling-up the number of images is also crucial to the maximum likelihood estimation of models that predict the image label, but



for which there are fewer than 10,000 examples per class. More recent generative pretraining or fine-tuning on large labeled and unlabeled datasets with dense or sparse labels has shown promise for a broader class of visual understanding tasks.

2. Background and Foundations

The principal imaging modalities leveraged by clinicians are computed tomography (CT), magnetic resonance imaging (MRI), X-ray imaging, ultrasound, and positron emission tomography (PET). Each modality has inherent strengths and weaknesses, making it particularly suited to the investigation of specific conditions. CT and MRI are the most sensitive for detecting subtle pathology in the thorax and abdomen; radiography enables rapid and cost-effective evaluation of the chest; ultrasound excels at documenting disease in the abdomen and obstetric applications; while PET is fundamentally different in probing metabolic activity rather than anatomical structure. Despite their contributions, for more complex cases, imaging results remain challenging to interpret, with inconsistent accuracies for specific conditions. Clinicians are, therefore, receptive to deploying AI-based diagnostic support tools to enhance performance. Indeed, recent work has shown deep learning models match or exceed radiologists when diagnostic performance is pooled across multiple diseases.

In supervised learning, convolutional neural networks (CNNs) have demonstrated clinical utility through supervised training on large-scale datasets. Subsequently, supervised pretraining on such datasets, followed by fine-tuning for specific diseases with limited data, has become a popular strategy in medical imaging diagnostics. A critical yet often overlooked consideration is the scaling of multimodal, multiclass datasets, because the training of CNNs requires hundreds of thousands to millions of labeled examples for optimal generalization.

Equation 1: 2D convolution (single channel $\rightarrow$ single channel)

Let the input image be $X \in \mathbb{R}^{H \times W}$, and a filter/kernel be $K \in \mathbb{R}^{k_h \times k_w}$. With stride s and padding p , the (cross-correlation style) convolution output is:

$$Y[i, j] = \sum_{u=0}^{k_h-1} \sum_{v=0}^{k_w-1} K[u, v] X[i \cdot s + u - p, j \cdot s + v - p]$$

Output spatial size

$$H_{\text{out}} = \left\lfloor \frac{H + 2p - k_h}{s} \right\rfloor + 1, W_{\text{out}} = \left\lfloor \frac{W + 2p - k_w}{s} \right\rfloor + 1$$

2.1. Medical Imaging Modalities

CT, MRI, X-ray, ultrasound, and PET are in widespread clinical use but are not equally represented within large data sources. Machine-learning diagnostic solutions should ideally be tested on imaging modalities from multiple systems; sources including the UK Biobank and others address this need. Nevertheless, the underlying principles are not necessarily applicable across



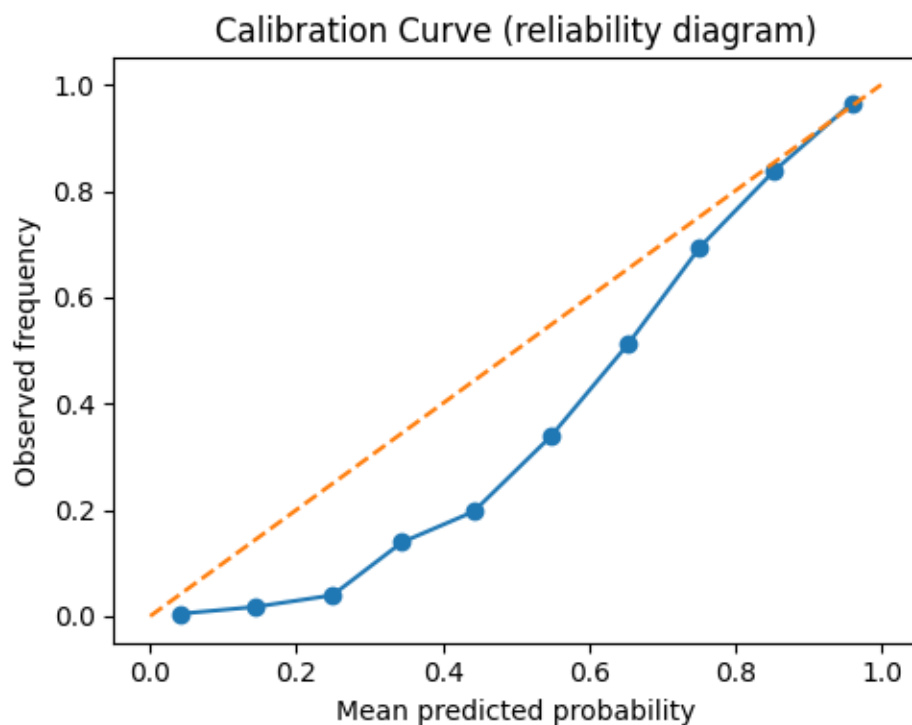
imaging modalities. Such concepts should therefore be described separately for each major imaging type, along with an exposition of their advantages and disadvantages.

As a primary screening tool, X-ray is effectively employed in certain clinical contexts. A number of algorithms for X-ray detection of abnormalities such as pneumonia, tuberculosis, and other lung pathologies have been reported. Potential pitfalls associated with bias and limited performance on unseen test sets are now receiving increasing attention. Moreover, there is a relatively limited pool of publicly available, annotated chest X-ray datasets. Some works have relied on abnormality-class predictive models to establish images as either normal or abnormal, while others have developed a continuum of disease-stage classifiers. Multiple strategies are available for detecting above-normal images, but recent developments emphasize the need for specialized models over one-size-fits-all approaches.

2.2. Deep Learning Principles for Diagnostics

Widespread application of deep learning in medical imaging diagnostics has explored Convolutional Neural Networks (CNNs) for supervised classification and detection. CNNs differ from prior deep-learning architectures in that they incorporate shared weights in filters, enabling feature learning from labeled data in a single end-to-end optimization. CNNs thus leverage deep structures with many levels to stepwise learn increasingly complex features. The input layer consists of 2D or 3D images, with small receptive fields centered on the cross-section in 3D structures. Task-specific output modules can convert CNN features into classification, detection, or segmentation results. Transfer learning has enabled the use of CNNs pretrained with large, publicly available, natural-imagery datasets to overcome common-label scarcity. Diagnostic-tasks-specific training of high-performance CNNs surpassed radiologist performance in detecting several thoracic diseases from frontal chest X-rays.

Other forms of label information, such as image-level distributions, may mitigate label-scarcity and support large-scale task-independent feature learning with CNNs. Large-scale and naturally formed multiclass label distributions in histopathological-image repositories have enabled true-label-scarcity alleviation via self-supervision, and CNN models pretrained on this naturally formed distribution outperformed supervised models. Deploying models in actual settings entails ensuring generalization to unseen data groups spanning origin datasets that vary in form or acquisition settings. Several adaptation approaches, primarily based on self-supervised and contrastive learning, have shown promise in improving generalization capacity. Probabilistic label modeling integrates label uncertainty into deep networks for joint label noise-tolerant classification and uncertainty estimation.



3. Data Acquisition and Curation

3.1. Sources of Large-Scale Clinical Data: Multicenter repositories, hospital data pipelines, and consortium-curated data have been increasingly made available for deep-learning research in many medical imaging applications. Given their large-scale nature, heterogeneity, and often combined open-source and clinical value, datasets from the Cancer Imaging Archive (TCIA) or the National Institutes of Health (NIH) represent very interesting opportunities to investigate the effectiveness of deep-learning algorithms at scale. Such data go beyond the limitations of a relatively small number of annotated images—an unsolved challenge for most imaging modalities and diseases. Similarly, the widely cited ImageNet dataset has opened a new direction for image classification via transfer learning, and the recent release of the UK Biobank MRI data is expected to stimulate new developments for the unsupervised or semi-supervised learning of organ atlases. These large, annotated training and testing data repositories, however, pose a different challenge: deep-learning models must demonstrate high generalization capabilities to unseen domains if they are to be successfully deployed in real-world clinical applications. Real-world clinical settings are often subject to scanner and vendor-related variation due to the complexity of imaging research, including the type of scanner and technical parameters such as slice thickness, contrast, and field strength. As clinical guidelines often favour a specific institution as the referring centre during the pre- and postprocessing pipelines, external validation is extremely difficult, if not impossible. Effectively tackling these domain shifts has thus become one of the key performance indicators for clinical imaging research. To

date, very few imaging-layer methods have been validated externally on datasets from a different institution, despite some research demonstrating the advantage of such evaluation protocols. Scalability to large-disparity multicenter studies remains an elusive goal, however, given the acknowledged low quality of public data and the longer data-acquisition pipelines for deep-learning-based image reconstruction or enhancement.

3.2. Data Quality, Labeling, and Annotation Protocols: The recent and widespread adoption of imaging technologies has made it possible to acquire a large amount of clinical data. However, generating these datasets requires special care and involves multiple layers of rules related to insurance, reimbursement, confidentiality, and ownership; especially in institutions where data privacy is strictly enforced. To address these challenges, a multilayer hierarchical labelling protocol has been designed for 15 diseases or conditions in patients diagnosed using different imaging modalities: computed tomography (CT) scans, magnetic resonance imaging (MRI) scans, X-rays, and autopsy cohort records. The labelling system adheres closely to the disease ontology system and attempts to serve as a dictionary of disease coding. The core goal is to build a set of labels that can be inherited or adapted in training shared models using the same, similar, or even different imaging pipelines, thereby forming a multi-institutional, multinational, multimodal, and multimajor-disease dataset for deep-learning-based imaging diagnostics. The labelling standard covers the full cycle of data management from acquisition to presentation of deep-learning models: generating ground-truth labels for supervised learning, annotating labels using medical-expert knowledge for implementation in the imaging layer, establishing VOI sampling and image-merging protocols during quality-check procedures, and estimating the hidden-node scores of supervised/classification deep-learning networks for clinical usage. Furthermore, the effect of human-machine co-learning on the gradual evolution of the prediction/classification model has been preliminarily explored. Finally, a pipeline to expand the data pool with deep-learning-generation quality-check protocols and quality-control assessments has been proposed.





Fig 2: Data Acquisition and Curation

3.1. Sources of Large-Scale Clinical Data

Although data-from-multiple-centers diagnostic studies have been performed to a limited extent for isolated modalities such as CT and X-ray, no approach addressing key pitfalls and supported by extensive data diversity has yet been executed. Several large-scale clinical imaging datasets have begun to emerge from multicenter repositories and hospital consortiums. Beyond diagnostic accuracy, key factors include robustness under dataset shift, the ability to capture the full range of conditions the model is to distinguish, and the non-redundancy of the classes in the labelled data. The novelty of soft labeling for the creation of ground truth provides additional training signal and improves diagnostic performance.

Multicenter Quarterly Quality Control Committees have been established to ensure a systematic quality control-driven approach for all imaging modalities, and imaging data pipelines provide essential infrastructure for a wide range of future projects. Data from The Cancer Imaging Archive (TCIA) and GlaS2 (GlaS Segmentation Challenge 2017) have already been utilized, and other major resources such as the UK Biobank study and other chest and brain CT imaging databases are in the aggregation stage. The establishment of the national Hong Kong Medical Data Consortium and Grand-A data pipeline has further enhanced data diversity and scale, with additional data consortiums for pharmacy claim data and associated nuclear medicine coverage also in preparation.

Equation 2: Multi-channel convolution (typical CNN layer)

Input $X \in \mathbb{R}^{C_{in} \times H \times W}$. Each output channel c has its own kernel bank $K_c \in \mathbb{R}^{C_{in} \times k_h \times k_w}$ and bias b_c :

$$Y_c[i, j] = b_c + \sum_{m=1}^{C_{in}} \sum_{u=0}^{k_h-1} \sum_{v=0}^{k_w-1} K_{c,m}[u, v] X_m[i \cdot s + u - p, j \cdot s + v - p]$$

3.2. Data Quality, Labeling, and Annotation Protocols

High-quality labels ensure the accuracy and reliability of deep learning models. High-quality labels are obtained through protocols for labeling, supervision, and quality assurance. Labeling protocols specify whether data are natively labeled (CT, MRI, X-ray) or require supervision and label creation (US, PET). Labelers and annotators must be domain experts who complete the labeling task either natively or by defining the ground truth. Inter-rater reliability checks ensure consistent annotations. Labelers receive feedback from experts to increase consistency. Biased or uncertain labels undergo further labeling rounds. Each dataset is manually inspected by experts to remove low-quality samples. For privacy reasons, certain datasets undergo preprocessing steps to ensure that accurate data labels remain while removing private information. Such pipeline strategies can enhance the bias characteristics of data while reducing the effect of random noise and improving the generalization ability of the final model.

Quality assurance involves the preservation, confirmation, and joint curation of large-scale multi-site medical imaging data. The multi-site nature of large-scale clinical imaging datasets leads to significant variations among scans, making data quality control an important aspect of deep learning for medical imaging. A pooled quality-control database and a quality-assurance pipeline for imaging datasets accompany an FDA-cleared clinical pipeline for deep-learning-based diagnosis of pneumonia on chest X-rays.



A clinical database that generates diagnosis-grade labels directly from routine clinical decisions is also created. Quality control and assurance procedures verify the collected data, provide quality-control labels for specific data modalities, and incorporate high-quality multi-site data and enhanced clinical imaging datasets into the training pipeline of a deep-learning-based clinical diagnostic service.

4. Model Architectures and Training Strategies

Convolutional neural networks (CNNs) are dominant in deep learning applications for imaging data. CNNs leverage gating mechanisms to share parameters across the receptive field, allowing feature representation with fewer training parameters compared to fully connected architectures. A convolutional layer consists of a set of filters that modulate the information in a localized patch of the input at the same spatial location, serving as an implicit spatiotemporal pooling and feature distraction layer. Individual filter responses (also called channels) within each feature map are combined into a single sum. High-level features are learned through multiple convolutional layers, with pooling layers used to reduce the spatial resolution of the feature maps and control the aspect ratio of the features.

To combine the advantages of unsupervised representation learning and supervised classification, pretraining is commonly employed, using models pretrained on large-scale datasets such as ImageNet or medical imaging repositories. Pretraining serves two purposes. The filters are adapted to initial feature extraction layers so that low-level features are learned more efficiently, and high-level features are trained with a large set of samples. Second, features at all layers are likely to be better performing compared to without pretraining because of the shared weights across a large number of samples. Data augmentation is also useful and widely applied in image classification tasks. Specific augmentations such as RecNet, mixup, and CutMix can further improve generalization performance.

More recently, vision transformers based on self-attention networks have demonstrated promising results on large-scale image classification tasks. Input images are divided into a sequence of patches that are fed into the transformer architecture, and the self-attention mechanism enables the model learning the similarity between different patches. Vision transformers avoid inductive biases from convolution and pooling operations. In addition, multimodal fusion transformers can enable joint modeling of the image input and tabular data such as electronic health record data. Nonetheless, transformers typically exhibit higher costs due to the quadratic computational and memory complexities with respect to the input sequence length. Efficient vision transformers can mitigate this cost by reducing the input dimensionality or incorporating convolutional layers within the model.

Equation 3: Backpropagation through convolution (core gradients)

Assume a loss L depends on Y . Let upstream gradient be $G[i, j] = \frac{\partial L}{\partial Y[i, j]}$.

Gradient w.r.t. kernel

$$\frac{\partial L}{\partial K[u, v]} = \sum_{i, j} G[i, j] X[i \cdot s + u - p, j \cdot s + v - p]$$

Gradient w.r.t. input

$$\frac{\partial L}{\partial X[a, b]} = \sum_{i, j} \sum_{u, v} G[i, j] K[u, v] \mathbf{1}\{a = i \cdot s + u - p, b = j \cdot s + v - p\}$$

4.1. Convolutional Neural Networks for Imaging

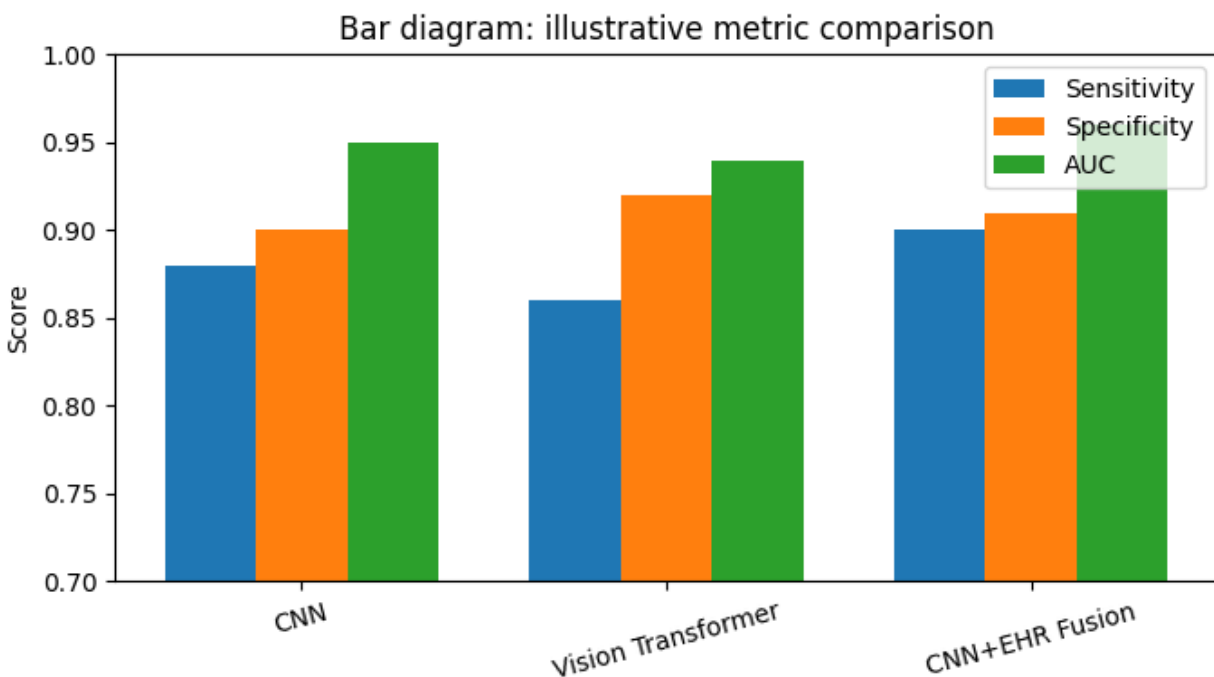
Convolutional Neural Networks (CNNs) are widely adopted for a variety of imaging tasks due to their success in large-scale natural-image classification. CNNs are designed to extract features at many levels of abstraction while using a structure that minimizes the number of parameters. By having multiple processing layers, CNNs apply backpropagation to update the weights and biases of several neurons, which are connected to the next layer neurons—and thus share the same weights. In image-centric applications, convolutional neural network architectures are often constructed in depth. Depth is introduced by stacking multiple convolutional layers followed by peeling-back or pooling layers. Each pooling layer down-samples the input, while preceding convolutional layers detect low-level features. In the initial layers, the receptive fields of the directional detectors are small enough to distinguish textures, forms, and edges. With a sufficiently large depth, the final layers aggregate these low-level components into complex and high-level objects or patterns. In medical applications, augmentations such as affine transformations, morphological filtering, noise, contrast, and flipping are commonly employed during the training process.

Pretraining on a large dataset and reusing the weights on the target task with transfer learning reduces the need for data in the training stage and alleviates the problem of overfitting. As in natural-image datasets with thousands of classes, transfer learning can also be applied in medical data. Tetrad systems and combinations of different modalities are validated many times for COVID-19 diagnosis by integrating X-ray, Computed Tomography, Computed Tomography angiograms, and others. Transfer learning is being used even for threat prediction using X-ray images. With scalable and diverse datasets in different modalities, CNNs are also being validated in the deployment phase, transferring directly into alimentary organ diagnosis using small-target endoscopic images, and are being explored as a method for detecting several rare diseases in large-scale data.

4.2. Transformer-Based Approaches in Medical Imaging

Recently, transformers have begun to achieve state-of-the-art performance in a variety of domains, through architectures such as Attention Is All You Need for natural language processing and Vision Transformers for image classification. Its self-attention mechanism processes all parts of the input simultaneously while yielding a global field of view at every stage. These advantages offer new opportunities for efficiency, multi-modal fusion, and development in the medical imaging domain. The output of the self-attention module can provide interpretable gradients for visualization by backpropagating through the calculated attention maps.

The key idea of the Vision Transformer is to divide images into small patches and project these patches linearly before feeding them into a transformer encoder. The resulting sequences of patch embeddings can thus be treated as words in a sentence. Vision Transformers leverage positional embeddings to preserve spatial information. Pretraining on large datasets in a self-supervised manner demonstrates the capability of Vision Transformers and enables fine-tuning for image classification with limited labels. Vision Transformer-based models can also be effectively used for medical image analysis.



5. Evaluation Metrics and Validation Frameworks

Evaluating deep learning models requires a thorough validation framework that encompasses diagnostic accuracy, robustness, calibration, uncertainty, and interpretability. Diagnostic accuracy is assessed via conventional metrics, including sensitivity, specificity, and area under the receiver operating characteristic curve (AUC). Robustness is evaluated through performance on a held-out test set, stratified K-fold cross-validation, and external validation using unseen data from multiple hospitals. Accepted metrics for each case and disease domain are reported in respective publications. Calibration and uncertainty are evaluated using reliability diagrams and appropriate methods for confidence estimation. Explanatory techniques and clinical assessments are deployed to examine model interpretability.

For classification tasks, sensitivity, specificity, and the associated confusion matrix are representatives of model performance but provide little insight into execution reliability or correctness. Therefore, sensitivity, specificity, and area under the receiver operating characteristic curve were complemented by external validation on institutionally segregated data from different imaging centers and by stratified K-fold cross-validation. External validation on real-world test datasets is an essential indicator of a generalizable machine-learning model's practical value. High-performance investments for creating deep learning classifiers using initial multicenter image data repositories earn further returns if those investments translate into improved classification performance on truly independent external test datasets.



5.1. Diagnostic Accuracy and Robustness

Diagnostic performance is often quantified with evaluation metrics of accuracy. In binary classification tasks, sensitivity, specificity, and area-under-the-curve (AUC) are routinely used as performance indicators. Sensitivity measures the proportion of true disease cases that are correctly classified, while specificity quantifies the ability of the model to identify healthy cases. AUC reports the average true positive rate across all possible thresholds. However, these metrics do not reflect the actual diagnostic performance of the model when applied in a clinical setting. A calibration curve, plotting true positive rates against the average confidence of predictions for each bin, reveals how well the predicted probabilities correspond to actual probabilities. In practice, components of the research literature, site-specific clinical models, or pipelines may not be subject to external validation. The validity of these components can be approximately assessed via existence of a hold-out test set, using n-fold cross-validation, or substituting a source site model into another site's pipeline.

Diagnostic models should be robust to small perturbations in performance. A simple approach is to check for changes when images are slightly transformed. For image classifications, the test set can be augmented during evaluation by applying transformations to each test image and averaging performance across the augmented set. Some perturbations—image blurriness, label-preserving rotations—can directly be induced during testing. A more systematic analysis makes use of simulated noise that mimics image acquisition variability. These methodologies are particularly useful for transformer-based models. The increased mobility of tokenization introduces additional implementation choices that radically affect the model's performance; therefore, an on-the-fly analysis of model reliability, risk, and failure characteristics becomes crucial to its clinical adoption.

Fig 3: Diagnostic Accuracy and Robustness

5.2. Calibration, Uncertainty, and Interpretability

Two aspects of DL model performance are frequently overlooked in the medical literature: calibration of predicted probabilities and predictive uncertainty. Calibration describes a relationship between predicted and observed frequencies—for example, if a



model outputs a predicted probability of 0.7 for a given class, the model should be correct for approximately 70 out of 100 instances with that same predicted probability. Calibration curves, which plot observed frequency against predicted probability, can visually summarize a model's calibration quality. Accurate calibration helps a clinician understand how much to trust a model's predictions, while poor calibration indicates that predicted probabilities cannot be reliably interpreted, limiting their usefulness in clinical decision making.

Beyond calibration, predictive uncertainty quantifies how much a model knows about its predictions. A well-calibrated and reliable model can predict confidently when it is sure and hedge its bets when it has little or no relevant information. Epsilon tokenization, layer-wise adaptive rate scaling, simple Monte-Carlo dropout, and temperature scaling are just a few of the possible methods for estimating uncertainty during inference. The information gleaned from translated uncertainty estimates maybe even more useful than the model predictions themselves. If a model could indicate that it believes it has unseen examples from a different distribution, then the key patient data could be excluded from decision making and modern clinical governance could apply during a subsequent clinician-led consult.

As AI-based clinical decision-support systems enter routine workflows, it is critical to ensure that the predictive risk for complex, real-world tasks is both correctly calibrated and reliably estimated. Consideration of these two aspects enables safer integration of any AI models into the clinical setting. Similarly, interpretability can be enhanced through a variety of techniques—starting with visualizing activation maps—and recent user studies have shown that model explainability increases clinician confidence in AI-assisted predictions. Finally, understanding model performance across diverse patient subgroups is essential for developing trustworthy AI tools. For example, bias in facial recognition systems led to markedly different accuracy for Black and White individuals.

Equation 4: Classification head equations (probabilities + loss)

Let the CNN/Transformer produce a feature vector $h \in \mathbb{R}^d$. Logit:

$$z = w^T h + b$$

Probability (sigmoid):

$$\hat{p} = \sigma(z) = \frac{1}{1 + e^{-z}}$$

For label $y \in \{0,1\}$:

$$L = -[y \log(\hat{p}) + (1 - y) \log(1 - \hat{p})]$$

Derivative w.r.t. logit

1. $\frac{\partial L}{\partial \hat{p}} = -\left(\frac{y}{\hat{p}} - \frac{1-y}{1-\hat{p}}\right)$
2. $\frac{\partial \hat{p}}{\partial z} = \hat{p}(1 - \hat{p})$



3. Multiply:

$$\frac{\partial L}{\partial z} = \hat{p} - y$$

6. Deployment in Clinical Settings

Embedding diagnostic algorithms into clinical workflows and practices can facilitate their real-world utility. Usually, their inclusion occurs at scan acquisition or diagnosis stages. They can help radiologists prioritize exams, offer second opinions, increase diagnostic accuracy, and assist in detecting abnormalities less encountered through training. Algorithms can also complement radiologist training, guiding less experienced professionals in specific areas or subspecialties.

Consequently, user interfaces that streamlined interaction between the diagnostic tool and the user played an essential role in development. Decisions on alerting thresholds, active steering toward abnormal regions, and volume-signaling interfaces influencing response times were particularly key. Alarm fatigue, either through overt notifications or by delayed warnings, is one possible consequence of the associated workload. The majority of interfaces inform the user about videoconferencing, online gaming, or other task elements, and fail to reduce time to detection in these situations.

Equation 5: Vision Transformer equations (patches + self-attention)

Image $X \in \mathbb{R}^{H \times W \times C}$. Patch size $P \times P$. Number of patches:

$$N = \frac{H}{P} \cdot \frac{W}{P}$$

Flatten each patch to $x_i \in \mathbb{R}^{P^2 C}$. Linear projection to model dimension D :

$$e_i = x_i E \text{ where } E \in \mathbb{R}^{(P^2 C) \times D}$$

Add positional embedding $p_i \in \mathbb{R}^D$:

$$t_i = e_i + p_i$$

Given token matrix $T \in \mathbb{R}^{N \times D}$, compute:

$$Q = TW_Q, K = TW_K, V = TW_V$$

with $W_Q, W_K, W_V \in \mathbb{R}^{D \times d_k}$.

Attention weights:



$$A = \text{softmax} \left(\frac{QK^T}{\sqrt{d_k}} \right)$$

Output:

$$\text{Attn}(T) = AV$$

For heads $h = 1..H$:

$$\text{MHA}(T) = \text{Concat}(\text{Attn}_1(T), \dots, \text{Attn}_H(T))W_O$$

6.1. Workflow Integration and User Interfaces

Based on established workflows within a medical institution, integration points for deep learning algorithms must be identified and an interface developed that enables interaction suitable for clinical use. Deep learning algorithms are usually conceived as stand-alone software packages that offer prediction with a graphical representation. Although this enables intuitive interaction, deployment in a clinical setting requires integration into existing medical workflows, often with the model functioning as a decision-support system for the clinician. In such cases, the clinician decides the final diagnosis, whereas the model serves only to strengthen or back up the diagnosis established by the clinician. When the model agrees with the diagnosis established by the clinician, the justifiable reliance on the model is enhanced. If, however, the model diagnosis contradicts the judgment of the clinician, this could indicate a problem requiring further investigation.

Thus, interaction with the model is more complex than merely receiving the prediction with its associated confidence value. Flowcharts representing the integration of deep learning into CT lung cancer diagnosis or CT kidney stone-related complications have been proposed. In these studies, the clinician presents the image to the model and receives assistance in diagnosing potential lung cancer or kidney stones. Point and click are the only necessary interactions from the clinician; the model under the hood does the entire processing in a way that is entirely transparent to the user. A similar principle can be established for integration with a Web application that serves as a graphic user interface (GUI) for models handling multi-class diagnosis. This requires integration of all models within a single application so the clinician can browse through the GUI and feed images for prediction to the relevant model while receiving the classification result and confidence.

The principal components of such a GUI are the user registration, serving of predictions through the model, monitoring of hard and easy samples, and intake of images for which predictions differ from the user's diagnosis. These components enable the GUI user to register as a new clinician, interact with the models, discover hard samples, consult with data scientists in case of model failure on certain samples, and create a testbed of hard images suitable for model retraining. Thus, all the functions necessary to furnish the model with information useful for further improvements are surfaced for clinical use in a simple manner.

6.2. Regulatory Compliance and Safety

Medical devices and software used in clinical environments must satisfy stringent regulatory requirements not only concerning technical performance and clinical utility but also safety. This entails proving that any harm attributable to the system — whether direct or collateral, immediate or delayed, during normal or anomalous operation — is negligible in relation to the clinical benefit



with respect to targeted patient cohorts. Such requirements have grown more complex and stringent over time, owing in part to software's extending role and to public awareness of adverse events, and institutions are increasingly applying risk management strategies. Consequently, safety testing and risk management are becoming integral elements of many product-development strategies. Combined with the own novelty of ML algorithms, this may hamper their integration into clinical practice.

Research groups undertaking the development of ML-based clinical-support tools should ensure that they conform to applicable medical-device regulations as early as possible, and ideally approach industrial teams early in the project. During the product's technical-development phase, usability (the user experience when interacting with a product) and HMI development often take place iteratively under a user-centered-design strategy, one that employs a user-experience strategy. This constitutes the main component of medical-device software development and encompasses the design of graphical user interfaces to communicate decision-support information to healthcare professionals, whether automated, semi-automated, or fully manual.

7. Ethical, Legal, and Social Implications

Medical imaging algorithms are developed and evaluated on publicly available datasets or large-scale clinical datasets that have been approved for research use. Nevertheless, there remain ethical, legal, and social implications that must be considered. These implications can arise from the acquisition, use, and interpretation of imaging data, alongside the built algorithms.

Bias can appear in data due to differences in the way different subgroups present, receive treatment, and are diagnosed. Strong bias may impede the algorithm's generalization abilities and could lead to misdiagnoses for underrepresented groups. Fairness is the extent to which the diagnostic accuracies across various subgroups are comparable. Fairness issues may be exacerbated if the model is directly deployed in a decision-making role. Equity relates to the access to, and deployment of, the model, which ideally should not be limited to a few countries and patient populations. Addressing bias, fairness, and equity requires representative data, ongoing evaluation, and targeted solutions, including strategies that help guarantee balanced performance across subgroups with known disparities.

7.1. Bias, Fairness, and Equity

Important diagnostic tasks must ensure both validity and fairness across population subgroups. Disparities in disease prevalence, risk, and imaging characteristics can emerge, as can performance gaps between groups. Sampled populations must therefore be representative of the task's intended demographic, including vulnerable subgroups, to avoid inequitable health outcomes.

Inherent biases in pretrained models can propagate even with full-factorial cross-validation and external assessment. Additional validation on minority groups becomes essential, while fairness metrics—such as disparate impact and equal opportunity—quantify nondiscrimination. Reject option classifier ensembles can mitigate unbalanced performance across subgroups. Tensor decomposition and batch normalization techniques further limit subgroup bias and enable tests on unseen subpopulations. Finally, diversifying the training data, e.g., through synthetic data generation, promotes accurate models for the underrepresented group and helps guarantee diagnostic performance is equitable.



Fig 4: Bias, Fairness, and Equity

7.2. Privacy Preservation and Data Sovereignty

Protection of healthcare data is paramount, and such data must maintain their sovereignty. Privacy-preserving technologies, including de-identification and legislation, can enable sharing data while safeguarding individual privacy. The diagnosis of medical conditions using medical imaging data (such as MRI, CT, and X-ray images) is an important activity of the healthcare industry. Artificial Intelligence provides essential techniques for automating such tasks; however, AI techniques often require large amounts of training data for generalization. Yet, in practice, acquiring a large amount of data can be difficult due to the high cost of obtaining reliable data with corresponding labels. AI systems for medical imaging can therefore become non-generalizable.

A natural solution is to use real medical imaging data from multiple hospitals across multiple countries. Although such datasets can be diverse and enormous, their use is hindered by patient privacy concerns. De-identification of patient data is a way to reduce the risk of privacy leakage and allow this sharing of data. Such de-identification can mechanically remove patient name, date of birth, and other patient identifiers easily. However, the radiologist's report is very critical for the diagnosis process. For example, if a patient has commonly occurring features (e.g. with white skin), it would be very helpful for the AI application to know such information for better generalization. However, the report does not have this information since it is considered identifiable.



8. Conclusion

Deep learning models for medical image diagnostics, trained on large-scale clinical datasets, promise to narrow performance gaps, enhance human and machine synergy, and enable broader deployment and use. These studies offer valuable insights into clinical translation, a critical yet often overlooked aspect of artificial intelligence applications in health care. Nevertheless, future deep-learning approaches must expand beyond pathological categories for which large-scale datasets exist, take into account patients' sex, race, and age to foster fairness and equity, protect patient privacy, ensure compliance with data sovereignty laws, and maintain societal trust.

The most significant advances in deep learning-based medical imaging diagnostics in recent years have resulted from the creation, sharing, and use of large-scale clinical datasets. It is anticipated that large-scale datasets will continue to empower large-scale studies that may be deployed in clinical settings or be translated into clinically relevant research questions. However, investigation of potential limitations—insufficient representativeness, fairness, and equity of the dataset, for instance—and appropriate mitigation strategies will remain essential.

8.1. Future Trends

Despite promising research progress, deploying deep learning diagnostic solutions at scale presents challenges. Integrating diagnostic tools into clinical workflows requires careful consideration of technical, regulatory, safety, and user-experience aspects. Moreover, verifying tool accuracy and robustness on local data before use is essential. Establishing comprehensive validation protocols remains paramount. Regulatory scrutiny is also growing to ensure that AI systems deployed in hospitals meet high standards. Developers must therefore pay appropriate attention to manage, assess, and reduce risks. Like other medical devices, AI-based diagnostic systems are subject to the same stringent standards as other healthcare devices.

Looking further ahead, as public and private institutional investments in AI research and development grow, exciting new developments are expected, especially in the fields of natural language processing, image synthesis, reinforcement learning, multimodal learning, and model efficiency. AI has the potential to drive significant advances in global healthcare and reshape many areas of medicine, including diagnostics, therapeutics, drug discovery, genomics, and basic science. However, achieving such a revolution is unlikely without an explicit focus on the ethical and practical challenges involved.

9. References

1. McKinney, S. M., Sieniek, M., Godbole, V., Godwin, J., Antropova, N., Ashrafian, H., Back, T., Chesus, M., Corrado, G. S., Darzi, A., Etemadi, M., Garcia-Vicente, F., Gilbert, F. J., Halling-Brown, M., Hassabis, D., Jansen, S., Karthikesalingam, A., Kundu, S., Ledsam, J. R., ... Shetty, S. (2020). International evaluation of an AI system for breast cancer screening. *Nature*, 577, 89–94.

2. Kim, H. E., Kim, H. H., Han, B. K., Kim, K. H., Han, K., Nam, H., Lee, E. H., & Kim, E. K. (2020). Changes in cancer detection and false-positive recall in mammography using artificial intelligence: A retrospective, multireader study. *The Lancet Digital Health*, 2(3), e138–e148.
3. Mashetty, S., Malempati, M., Paleti, S., Adusupalli, B., & Singireddy, J. (2025). A Multidisciplinary Framework for AI and Data-Driven Transformation in Taxation, Insurance, Mortgage Financing, and Financial Advisory: Integrating Cloud Computing, Deep Learning, and Agentic AI for Community-Centric Economic Development. *Insurance, Mortgage Financing, and Financial Advisory: Integrating Cloud Computing, Deep Learning, and Agentic AI for Community-Centric Economic Development*.
4. Harmon, S. A., Sanford, T. H., Xu, S., Turkbey, E. B., Roth, H., Xu, Z., Yang, D., Myronenko, A., Xu, D., Turkbey, B., Wang, A., Pathak, S., Bodian, C. A., Elnakib, A., Xu, Z., Wang, D., Xu, Y., Freifeld, A. G., ... Choyke, P. L. (2020). Artificial intelligence for the detection of COVID-19 pneumonia on chest CT using multinational datasets. *Nature Medicine*, 26, 1605–1611.
5. Li, L., Qin, L., Xu, Z., Yin, Y., Wang, X., Kong, B., Bai, J., Lu, Y., Fang, Z., Song, Q., Cao, K., Liu, D., Wang, G., Wu, H., Han, B., & Chen, L. (2020). Using artificial intelligence to detect COVID-19 and community-acquired pneumonia based on pulmonary CT: Evaluation of the diagnostic accuracy. *Radiology*, 296(2), E65–E71.
6. Paleti, S., Baliyan, M., Aitha, A. R., Reddy, B. A., Bhadauria, G. S., & Sing, S. A. (2025, August). Graph—LSTM Hybrid Model for Improving Fraud Detection Accuracy in E-Commerce Financial Services. In *2025 2nd International Conference on Intelligent Algorithms for Computational Intelligence Systems (IACIS)* (pp. 1-6). IEEE.
7. Mei, X., Lee, H.-C., Diao, K.-y., Huang, M., Lin, B., Liu, C., Xie, Z., Ma, Y., Robson, P. M., Chung, M., Bernheim, A., Mani, V., Fuster, V., Morris, E. A., & Yang, Y. (2020). Artificial intelligence-enabled rapid diagnosis of patients with COVID-19. *Nature Medicine*, 26, 1224–1228.
8. Zhou, S. K., Greenspan, H., Davatzikos, C., Duncan, J. S., van Ginneken, B., Madabhushi, A., Prince, J. L., Rueckert, D., & Summers, R. M. (2021). A review of deep learning in medical imaging: Imaging traits, technology trends, case studies with progress highlights, and future promises. *Proceedings of the IEEE*, 109(5), 820–838.

9. Burugulla, J. K. R., Kannan, S., Malempati, M., Challa, H. A. H. S. R., Pandugula, C., & Goma, T. (2025, April). Federated Learning Based Cloud Computing Solutions with Privacy Preservation for Intelligent Utilities. In 2025 4th OPJU International Technology Conference (OTCON) on Smart Computing for Innovation and Advancement in Industry 5.0 (pp. 1-6). IEEE.
10. Aggarwal, R., Sounderajah, V., Martin, G., Ting, D. S. W., Karthikesalingam, A., King, D., Ashrafiyan, H., & Darzi, A. (2021). Diagnostic accuracy of deep learning in medical imaging: A systematic review and meta-analysis. *npj Digital Medicine*, 4, 65.
11. Liu, X., Faes, L., Kale, A. U., Wagner, S. K., Fu, D. J., Bruynseels, A., Mahendiran, T., Moraes, G., Shamdas, M., Kern, C., Ledsam, J. R., Schmid, M. K., Balaskas, K., Topol, E. J., Bachmann, L. M., Keane, P. A., & Denniston, A. K. (2021). A comparison of deep learning performance against health-care professionals in detecting diseases from medical imaging: A systematic review and meta-analysis. *The Lancet Digital Health*, 1(6), e271–e297.
12. Kummari, D. N., Singireddy, J., Sheelam, G. K., Nandan, B. P., Pandiri, L., Lakkarasu, P., & Dwaraka. (2025, August). Generative AI Models for Process Optimization in Semiconductor Wafer Design and Yield Prediction. In *International Conference on Artificial Intelligence: Theory and Applications* (pp. 220-233). Cham: Springer Nature Switzerland.
13. Seah, J. C. Y., Tang, C. H. M., Buchlak, Q. D., Holt, X. G., Wardman, J. B., Aimoldin, A., & Chan, W. P. (2021). Effect of a comprehensive deep-learning model on the accuracy of chest X-ray interpretation by radiologists: A retrospective, multireader multicase study. *The Lancet Digital Health*, 3(8), e496–e506.
14. Tjoa, E., & Guan, C. (2021). A survey on explainable artificial intelligence (XAI): Toward medical XAI. *IEEE Transactions on Neural Networks and Learning Systems*, 32(11), 4793–4813.
15. Adusupalli, B., Malempati, M., Paleti, S., Mashetty, S., & Singireddy, J. (2025). Integrated financial ecosystems: AI-driven innovations in taxation, insurance, mortgage analytics, and community investment through cloud, big data, and advanced data engineering. *Journal of Information Systems Engineering and Management*, 10, 1103-1117.
16. Kelly, C. J., Karthikesalingam, A., Suleyman, M., Corrado, G., & King, D. (2019). Key challenges for delivering clinical impact with artificial intelligence. *BMC Medicine*, 17, 195.



17. Ghassemi, M., Oakden-Rayner, L., & Beam, A. L. (2021). The false hope of current approaches to explainable artificial intelligence in health care. *The Lancet Digital Health*, 3(11), e745–e750.
18. Chen, X., Wang, X., Zhang, K., Fung, K.-M., Thai, T. C., Moore, K., Mannel, R. S., Liu, H., Zheng, B., & Qiu, Y. (2022). Recent advances and clinical applications of deep learning in medical image analysis. *Medical Image Analysis*, 79, 102444.
19. Rajpurkar, P., Chen, E., Banerjee, O., & Topol, E. J. (2022). AI in health and medicine. *Nature Medicine*, 28, 31–38.
20. Suura, S. R., Chava, K., Chakilam, C., Nuka, S. T., Maguluri, K. K., & Goma, T. (2025, April). Blockchain-Based Secure and Scalable Models for Healthcare Network Traffic Monitoring and Optimization. In *2025 4th OPJU International Technology Conference (OTCON) on Smart Computing for Innovation and Advancement in Industry 5.0* (pp. 1-6). IEEE.
21. Kelly, C. J., Young, A. J., Pearson, D., Ranganathan, P., & Karthikesalingam, A. (2022). Assessing the performance and generalisability of deep learning systems in medical imaging. *Medical Image Analysis*, 78, 102402.
22. Chalkidou, A., Shokraneh, F., & Kijauskaite, G. (2022). Recommendations for the development and use of imaging test sets to investigate the test performance of artificial intelligence in health screening. *The Lancet Digital Health*, 4(12), e899–e905.
23. Seyyed-Kalantari, L., Zhang, H., McDermott, M. B. A., Chen, I. Y., & Ghassemi, M. (2021). Underdiagnosis bias of artificial intelligence algorithms applied to chest radiographs in under-served patient populations. *Nature Medicine*, 27, 2176–2182.
24. Kalisetty, S., Lakkarasu, P., Singireddy, S., Burugulla, J. K. R., Challa, K., & Gadi, A. L. (2025, August). Next-Gen Payment Gateways: Leveraging Federated Learning for Fraud Detection in Cross-Border Transactions. In *International Conference on Artificial Intelligence: Theory and Applications* (pp. 200-211). Cham: Springer Nature Switzerland.
25. Yang, Y., Zhang, H., Gichoya, J. W., Katabi, D., & Ghassemi, M. (2024). The limits of fair medical imaging AI in real-world generalization. *Nature Medicine*, 30, 2838–2848.
26. Zhang, H., Dullerud, N., Seyyed-Kalantari, L., & Ghassemi, M. (2022). Algorithmic bias in medical imaging: A systematic review. *npj Digital Medicine*, 5, 145.

27. Oakden-Rayner, L., Gale, W., Bonham, T. A., & Lungren, M. P. (2020). Predicting patient outcomes from medical imaging using deep learning: A systematic review. *Medical Image Analysis*, 64, 101729.
28. Zhang, H., Li, J., Yang, X., & Wang, J. (2020). Deep learning for medical image analysis: A review. *IEEE Access*, 8, 124423–124438.
29. Ardila, D., Kiraly, A. P., Bharadwaj, S., Choi, B., Reicher, J. J., Peng, L., Tse, D., Etemadi, M., Ye, W., Corrado, G., Naidich, D. P., & Shetty, S. (2019). End-to-end lung cancer screening with three-dimensional deep learning on low-dose chest computed tomography. *Nature Medicine*, 25, 954–961.
30. Mashetty, S. (2025). Technology-driven analytics in mortgage-backed securities for single-family mortgage financing. Available at SSRN 5236573.
31. Kather, J. N., Heij, L. R., Grabsch, H. I., Loeffler, C., Echle, A., Muti, H. S., Krause, J., Niehues, J. M., Arvaniti, E., Haß, C., et al. (2020). Pan-cancer image-based detection of clinically actionable genetic alterations. *Nature Cancer*, 1, 789–799.
32. Bilal, M., Raza, S. E. A., Azam, A., Graham, S., Ilyas, M., Cree, I. A., & Snead, D. (2021). Development and validation of a weakly supervised deep learning framework to predict the molecular subtype of primary colorectal cancer from histology images. *Medical Image Analysis*, 70, 101996.
33. Campanella, G., Hanna, M. G., Geneslaw, L., Miraflor, A., Werneck Krauss Silva, V., Busam, K. J., Brogi, E., Reuter, V. E., Klimstra, D. S., & Fuchs, T. J. (2019). Clinical-grade computational pathology using weakly supervised deep learning on whole slide images. *Nature Medicine*, 25, 1301–1309.
34. Echle, A., Rindtorff, N. T., Brinker, T. J., Luedde, T., Pearson, A. T., & Kather, J. N. (2021). Deep learning in cancer pathology: A new generation of clinical biomarkers. *British Journal of Cancer*, 124, 686–696.
35. Sivanand, R., Kumar, D. P., Nagabhyru, K. C., Natarajan, E. P., Pamisetty, V., & Kapila, D. (2025, September). IoT and AI for Real-Time Monitoring in Substation Automation. In 2025 International Conference on Computing and Communications (COMPUTINGCON) (pp. 1-5). IEEE.
36. Bulten, W., Pinckaers, H., van Boven, H., Vink, R., de Bel, T., van Ginneken, B., van der Laak, J., Hulsbergen-van de Kaa, C., & Litjens, G. (2022). Automated deep-learning system

46. Prasad, V. K., Verma, A., Bhattacharya, P., Shah, S., Chowdhury, S., Bhavsar, M., Aslam, S., & Ashraf, N. (2024). Revolutionizing healthcare: A comparative insight into deep learning's role in medical imaging. *Scientific Reports*, 14, 30273.
47. Zhang, H., Qie, Y., & colleagues. (2023). Applying deep learning to medical imaging: A review. *Applied Sciences*, 13(18), 10521.
48. Moor, M., Banerjee, O., Abad, Z. S. H., Krumholz, H. M., Leskovec, J., Topol, E. J., & Rajpurkar, P. (2023). Foundation models for generalist medical artificial intelligence. *Nature*, 616, 259–265.
49. Pamisetty, V. (2019). Machine Learning Models for Real-Time Tax Fraud Detection and Risk Assessment in Digital Government Systems. *Global Research Development (GRD)* ISSN, 2455-5703.
50. Moor, M., Huang, Q., Wu, S., Yasunaga, M., Zakka, C., Dalmia, Y., Reis, E. P., et al. (2023). Med-Flamingo: A multimodal medical few-shot learner. *Proceedings of the 37th Conference on Neural Information Processing Systems*, 1–15.
51. Yang, Y., Hou, W., Zhang, H., Gichoya, J. W., & Ghassemi, M. (2023). A comprehensive study of deep learning models for medical imaging and their generalization across datasets. *Medical Image Analysis*, 85, 102749.
52. Izhar, A., Idris, N., & Japar, N. (2025). Medical radiology report generation: A systematic review of current deep learning methods, trends, and future directions. *Artificial Intelligence in Medicine*, 168, 103220.
53. Li, Y., Kong, C., Zhao, G., & Zhao, Z. (2025). Automatic radiology report generation with deep learning: A comprehensive review of methods and advances. *Artificial Intelligence Review*, 58, 344.
54. Laçi, H., Sevrani, K., & Iqbal, S. (2025). Deep learning approaches for classification tasks in medical X-ray, MRI, and ultrasound images: A scoping review. *BMC Medical Imaging*, 25, 156.
55. Ma, D., Pang, J., Gotway, M. B., & Liang, J. (2025). A fully open AI foundation model applied to chest radiography. *Nature*, 643, 488–498.
56. Yang, X., Chen, Y., & colleagues. (2025). Application of artificial intelligence in medical imaging: Current status and future directions. *iRADIOLOGY*.
57. Farhan, O. A. M., Alnaggar, F., Jagadale, B. N., Saif, M. A. N., Ghaleb, O. A. M., Ahmed, A. A. Q., Aqlan, H. A. A., & Al-Ariki, H. D. E. (2024). Efficient artificial intelligence



approaches for medical image processing in healthcare: Comprehensive review, taxonomy, and analysis. *Artificial Intelligence Review*, 57, 221.

58. Yang, Y., Zhang, H., Gichoya, J. W., Katabi, D., & Ghassemi, M. (2024). The limits of fair medical imaging AI in real-world generalization. *Nature Medicine*, 30, 2838–2848.
59. Cruz Rivera, S., Liu, X., Chan, A.-W., Denniston, A. K., Calvert, M. J., & SPIRIT-AI and CONSORT-AI Working Group. (2020). Guidelines for clinical trial protocols for interventions involving artificial intelligence: The SPIRIT-AI extension. *The Lancet Digital Health*, 2(10), e549–e560.