

Review

Not peer-reviewed version

Bridging the Gap Between Black Box AI and Clinical Practice: Advancing Explainable AI for Trust, Ethics, and Personalized Healthcare Diagnostics

[Dang Anh Tuan](#)*

Posted Date: 25 September 2024

doi: 10.20944/preprints202409.1974.v1

Keywords: Explainable AI (XAI); Healthcare diagnostics; Grad-CAM; SHAP; LIME; Trust in AI; Ethical AI; Bias mitigation; Regulatory compliance; Personalized care; AI interpretability; Clinical decision-making



Preprints.org is a free multidiscipline platform providing preprint service that is dedicated to making early versions of research outputs permanently available and citable. Preprints posted at Preprints.org appear in Web of Science, Crossref, Google Scholar, Scilit, Europe PMC.

Copyright: This is an open access article distributed under the Creative Commons Attribution License which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited.

Review

Bridging the Gap between Black Box AI and Clinical Practice: Advancing Explainable AI for Trust, Ethics, and Personalized Healthcare Diagnostics

Dang Anh Tuan

Faculty of Information Technology, Industrial University of Ho Chi Minh City, Ho Chi Minh City, 70000, Vietnam; tuanda222@gmail.com; Tel: +84987 066 222

Abstract: Explainable AI (XAI) has emerged as a pivotal tool in healthcare diagnostics, offering much-needed transparency and interpretability in complex AI models. XAI techniques, such as SHAP, Grad-CAM, and LIME, enable clinicians to understand AI-driven decisions, fostering greater trust and collaboration between human and machine in clinical settings. This review explores the key benefits of XAI in enhancing diagnostic accuracy, personalizing patient care, and ensuring compliance with regulatory standards. However, despite its advantages, XAI faces significant challenges, including balancing model accuracy with interpretability, scaling for real-time clinical use, and mitigating biases inherent in medical data. Ethical concerns, particularly surrounding fairness and accountability, are also discussed in relation to AI's growing role in healthcare. The review emphasizes the importance of developing hybrid models that combine high accuracy with improved interpretability and suggests that future research should focus on explainable-by-design systems, reducing computational costs, and addressing ethical issues. As AI continues to integrate into healthcare, XAI will play an essential role in ensuring that AI systems are transparent, accountable, and aligned with the ethical standards required in clinical practice.

Keywords: Explainable AI (XAI); Healthcare diagnostics; Grad-CAM; SHAP; LIME; Trust in AI; Ethical AI; Bias mitigation; Regulatory compliance; Personalized care; AI interpretability; Clinical decision-making

1. Introduction

1.1. Background and Motivation

Artificial Intelligence (AI) has become an indispensable tool in modern healthcare, revolutionizing various areas, particularly diagnostics. AI systems, especially machine learning (ML) and deep learning (DL) models, have shown remarkable success in automating the detection of diseases, predicting patient outcomes, and analyzing complex medical data, such as medical images and genomic sequences. For instance, convolutional neural networks (CNNs) have been widely used for detecting abnormalities in medical images, including tumor detection in MRI and CT scans, outperforming traditional diagnostic methods in both speed and accuracy [1]. Moreover, AI's ability to analyze big data has led to breakthroughs in personalized medicine and predictive analytics for patient outcomes, opening new avenues for precision healthcare [2].

However, the adoption of AI in clinical practice faces a critical challenge: the "black-box" nature of many AI models. While these models deliver highly accurate predictions, they often do so without providing clear explanations of how or why a particular decision was made. This opacity poses a significant issue in healthcare, where transparency is essential for clinical decision-making. Physicians need to understand AI's rationale to trust its recommendations and to communicate them effectively to patients. Moreover, healthcare decisions frequently carry life-or-death implications, and unexplained AI decisions could lead to a lack of accountability and increased legal risks [3].

This is where Explainable AI (XAI) comes into play. XAI seeks to provide clarity by making AI models more interpretable and transparent. The concept of XAI is crucial for bridging the gap between AI's decision-making processes and clinical needs, enabling healthcare professionals to understand and trust AI systems more fully. This is not just a technical challenge but also an ethical one, as explainability is fundamental for ensuring fairness and accountability in medical AI applications [4].

1.2. Scope and Objectives of the Review

The objective of this review is to provide a comprehensive analysis of the current state of XAI in healthcare diagnostics. We will examine the leading XAI techniques, including post-hoc explainability methods such as LIME and SHAP, as well as inherently interpretable models like decision trees and rule-based systems. The review will focus on how these methods are being applied to critical areas such as medical imaging, genomics, and predictive analytics for patient outcomes. Furthermore, we will explore the trade-offs between explainability and performance in AI models, discuss the ethical and regulatory implications, and identify the barriers to the widespread adoption of XAI in clinical practice.

In this review, we aim to address the key question: How can XAI improve the transparency and trustworthiness of AI systems in healthcare, and what are the remaining challenges to its full integration into medical practice? By investigating these issues, we hope to offer valuable insights for future research and development in this rapidly evolving field, ensuring that AI becomes not only an accurate but also an interpretable and reliable tool in healthcare diagnostics.

2. Overview of Artificial Intelligence in Healthcare Diagnostics

2.1. AI Techniques in Healthcare

AI has revolutionized healthcare diagnostics by introducing computational models that mimic human cognitive functions, primarily through ML and DL techniques. These methods have become indispensable in analyzing large, complex datasets, enabling precise diagnosis, prognosis, and treatment personalization. Among the most prominent AI techniques used in healthcare diagnostics are convolutional neural networks (CNNs), recurrent neural networks (RNNs), decision trees, and support vector machines (SVMs), each with its unique advantages and applications.

CNNs have become the gold standard for analyzing medical images due to their ability to automatically detect patterns and extract relevant features from data, often surpassing traditional image analysis methods. CNNs have been widely used in the detection of tumors, brain anomalies, and other pathologies from modalities such as magnetic resonance imaging (MRI), computed tomography (CT), and X-rays [1]. For instance, CNNs have been instrumental in automating breast cancer detection, significantly improving diagnostic accuracy and reducing human error in mammogram analysis [5].

RNNs and their variant, long short-term memory (LSTM) networks, have been particularly effective in processing sequential medical data, such as time-series data from electrocardiograms (ECGs) or patient health records. RNNs excel in capturing temporal dependencies, making them valuable for predicting disease progression or patient outcomes based on historical data [6]. LSTM models have been applied in predictive healthcare analytics, such as predicting heart failure readmissions or monitoring patients with chronic diseases.

In addition to deep learning techniques, traditional machine learning methods like decision trees and SVMs continue to play a role in healthcare diagnostics, especially in scenarios where interpretable models are required. Decision trees are favored for their simplicity and transparency, allowing healthcare professionals to understand the logic behind AI-driven decisions. SVMs, on the other hand, are powerful classifiers used for tasks such as gene expression analysis and early disease detection in oncology [7].

The most commonly used AI techniques in healthcare diagnostics are CNNs, RNNs, SVMs, and Decision Trees, each with specific applications and trade-offs. **Table 1** provides a comparative summary of these techniques in terms of their healthcare applications, advantages, and limitations.

Table 1. Comparison of AI Techniques in Healthcare Diagnostics.

Technique	Application in Healthcare	Advantages	Limitations	References
CNN	Medical imaging (e.g., MRI, CT, X-rays)	High accuracy, automatic feature extraction	Black-box nature, requires large datasets	[1]
RNN	Time-series data (e.g., ECGs, patient history)	Captures temporal dependencies in medical data	Difficult to train, prone to vanishing gradient	[6]
SVM	Cancer detection, gene expression analysis	Effective in small, high-dimensional datasets	Less effective for large, unstructured data	[7]
Decision Trees	Diagnosis support in clinical data	Easy to interpret, transparent decision-making	Prone to overfitting, lower accuracy compared to DL	[8]

Table 1 shows CNNs are primarily used in medical imaging due to their high accuracy in automatic feature extraction from large image datasets. However, their black-box nature makes them difficult to interpret [1]. RNNs are suitable for time-series data, such as patient histories or ECGs, but they are difficult to train and can suffer from vanishing gradient problems [6]. SVMs are highly effective in small, high-dimensional datasets, making them a strong choice for tasks like gene expression analysis, although they struggle with unstructured data [7]. Decision trees provide transparency and ease of interpretation but tend to overfit, especially in more complex healthcare datasets [8].

2.2. Challenges of Black-box AI in Healthcare

Despite the successes of AI in healthcare, the widespread adoption of AI systems in clinical diagnostics is hindered by several challenges, with the most prominent being the "black-box" nature of many advanced AI models, particularly deep learning. Black-box models, such as deep neural networks, are typically highly complex, comprising multiple layers of computation that make their decision-making processes opaque to users, including medical professionals. This lack of transparency poses significant risks in the healthcare domain, where clinical decisions must be interpretable, justifiable, and accountable.

In clinical practice, trust is a critical component of decision-making. Physicians and healthcare providers need to understand and justify the decisions suggested by AI models, especially in critical diagnoses where lives are at stake. The inability to explain how or why a model arrived at a particular diagnosis or prognosis can result in hesitancy among clinicians to adopt AI tools, even if these tools are more accurate than human judgment in some cases [9]. This lack of transparency not only undermines trust but can also have legal and ethical implications, particularly if a model's decision leads to an adverse outcome.

Moreover, healthcare is a highly regulated field, with stringent legal and ethical standards requiring full accountability for clinical decisions. When AI models provide no clear rationale for their outputs, it becomes challenging to validate or scrutinize their decisions in a legal or regulatory framework. For example, if an AI system recommends a particular treatment but fails to provide a comprehensible explanation for this recommendation, it could be difficult for clinicians to defend that decision in court if something goes wrong [10]. Furthermore, AI systems trained on biased or incomplete datasets can introduce unintended biases into medical decision-making, potentially leading to discriminatory outcomes, particularly for minority or underrepresented groups [11].

The inherent opacity of many AI models also complicates their integration into existing clinical workflows. Healthcare professionals often prefer models that provide interpretable and actionable insights, allowing them to incorporate AI recommendations into their decision-making processes confidently. However, black-box models typically offer little to no insight into their decision pathways, which contrasts with the need for clear, evidence-based reasoning in medicine.

In summary, while AI holds tremendous potential for revolutionizing healthcare diagnostics, its black-box nature remains a significant barrier to adoption. The need for interpretability, transparency, and accountability in AI-driven healthcare is paramount to ensuring that AI systems are trusted and integrated into clinical practice effectively. This challenge has given rise to the development of XAI, a set of methods aimed at making AI models more understandable and interpretable, which will be explored further in the next sections of this review.

Grad-CAM is a widely used method to provide visual explanations for CNNs in medical imaging. By applying Grad-CAM, clinicians can see the regions of medical images that the model has focused on when making a diagnostic decision [12]. As shown in **Figure 1**, the heatmap highlights the areas of importance in a brain MRI for detecting a tumor.

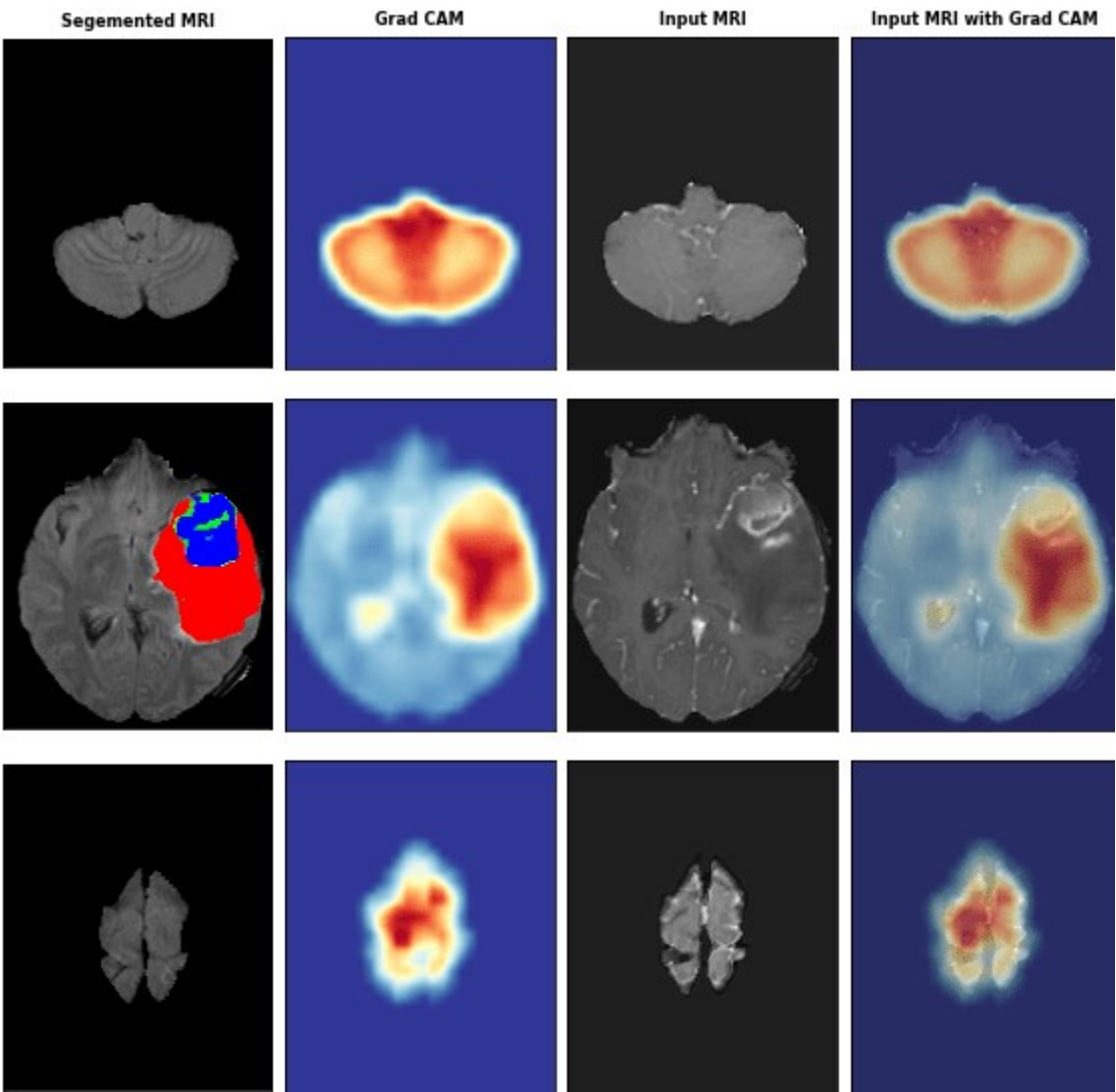


Figure 1. Grad-CAM visualizations generated at different depths of the same MRI volume. Row 2 highlights regions where the tumor exists, while Row 1 shows regions without the tumor from a bottom view, and Row 3 depicts regions without the tumor from a top view of the brain. The heatmap indicates areas of importance in the CNN's analysis, with warmer colors (red, orange) showing regions of high significance for tumor detection.

3. Explainable AI: Concepts and Techniques

3.1. Definition of Explainability

XAI is an evolving subfield within artificial intelligence that addresses the need to make machine learning models, particularly complex black-box models like deep neural networks, more transparent and interpretable. In the context of healthcare, where diagnostic decisions may have life-or-death consequences, explainability is critical for both ethical and practical reasons. Medical professionals require a clear understanding of how an AI model arrives at a specific diagnosis or prediction to ensure its reliability and to gain clinical trust in the system [4].

Explainability can be categorized into three main levels:

Model-Level Explainability: Provides a global view of how a model functions as a whole, such as decision trees or linear regression models, where decision paths and feature importances are inherently interpretable.

Decision-Level Explainability: Focuses on explaining specific predictions, such as why a model classified a particular scan as indicative of a tumor. Techniques like SHAP (SHapley Additive exPlanations) and LIME (Local Interpretable Model-agnostic Explanations) are used to achieve decision-level explanations [13].

Process-Level Explainability: Involves understanding how the AI model processes input data through different layers or components, a common practice in deep learning where visualization techniques like Grad-CAM (Gradient-weighted Class Activation Mapping) highlight relevant image regions that contribute to predictions [14].

With growing concerns about AI's "black-box" nature, XAI techniques have gained significant attention in healthcare diagnostics, where the ability to interpret and verify AI outputs is crucial for clinical validation and compliance with regulatory requirements.

3.2. Common XAI Techniques

XAI techniques fall broadly into two categories: post-hoc explainability methods, which are applied after a model has been trained, and inherently interpretable models, which are designed to be transparent from the start. In healthcare diagnostics, post-hoc methods are commonly used to explain predictions from high-performing but opaque models like CNNs.

Table 2. Comparison of Common Explainable AI Techniques in Healthcare.

Technique	Type	Explanation Process	Application in Healthcare	Limitations	References
LIME	Post-hoc	Locally perturbs input data to approximate the decision boundary of the model around the instance being explained.	Useful in explaining complex models for medical image classification or patient diagnosis.	Only provides local interpretability; less effective for very large datasets.	[15]
SHAP	Post-hoc	Based on cooperative game theory, assigns Shapley values to features that	Often applied in risk prediction models for chronic diseases, patient readmission, or	Computationally expensive for large models.	[13]

		contributed to a treatment prediction.	planning.		
Grab-CAM	Post-hoc	Highlights the regions of an image that are most relevant to the model's predictions using heatmaps.	Widely used in medical imaging for identifying regions of interest, such as tumors.	Limited to convolutional neural networks.	[14]
DeepLIFT	Post-hoc	Tracks the contribution of each input feature relative to a reference input, improving gradient-based methods.	Applied in genomics and precision medicine for feature attribution.	Less interpretable for complex temporal models.	[16]
Decision Trees	Inherently Interpretable	Visualizes decision-making through a tree of logical conditions, making predictions transparent.	Effective in rule-based diagnosis and decision support systems.	Prone to overfitting and less accurate in comparison to deep models.	[17]

Post-hoc Explainability Methods:

LIME: LIME works by perturbing the input data (e.g., changing pixels in an image) and observing how these changes impact the model's prediction. This method is widely used in interpreting image classification models and personalized medicine to explain individual diagnoses [15]. However, it is limited by its local nature, providing insight into a model’s behavior only around a specific prediction rather than globally.

SHAP: SHAP is grounded in cooperative game theory and calculates the Shapley value for each feature, offering a global perspective on feature importance across all predictions. SHAP is highly valuable in healthcare applications like risk prediction for chronic diseases or feature attribution in genomic studies [13]. The downside is its computational complexity, especially for deep learning models.

Grad-CAM: Grad-CAM is particularly useful in medical imaging, where it overlays heatmaps on input images to visualize which regions most influence the CNN’s decision. For example, it can highlight tumor regions in MRI scans, making the CNN’s reasoning transparent [14].

DeepLIFT: This method is an improvement on gradient-based techniques, calculating the contribution of each input feature relative to a baseline reference. It has been particularly useful in genomics, where it helps explain the influence of specific genes or sequences on a model’s output [16].

Inherently Interpretable Models:

Decision Trees: Decision trees remain popular in healthcare diagnostics for their simplicity and transparency, where each decision is based on clear, logical steps. They are used in decision support systems, such as diagnosing diseases based on clinical guidelines. However, decision trees often suffer from overfitting and lower accuracy compared to more complex models like deep neural networks [17].

The **Table 2** compares different XAI techniques, illustrating their respective applications in healthcare diagnostics. In particular, post-hoc methods like SHAP and LIME provide detailed explanations of individual predictions, while Grad-CAM is more suited to image-based tasks. **Figure 2** demonstrates SHAP values in a risk prediction model for chronic diseases, illustrating how each feature contributes to the final prediction. This type of visual explanation is crucial for understanding how an AI model evaluates various risk factors, making the model’s decision-making process more transparent to healthcare professionals.

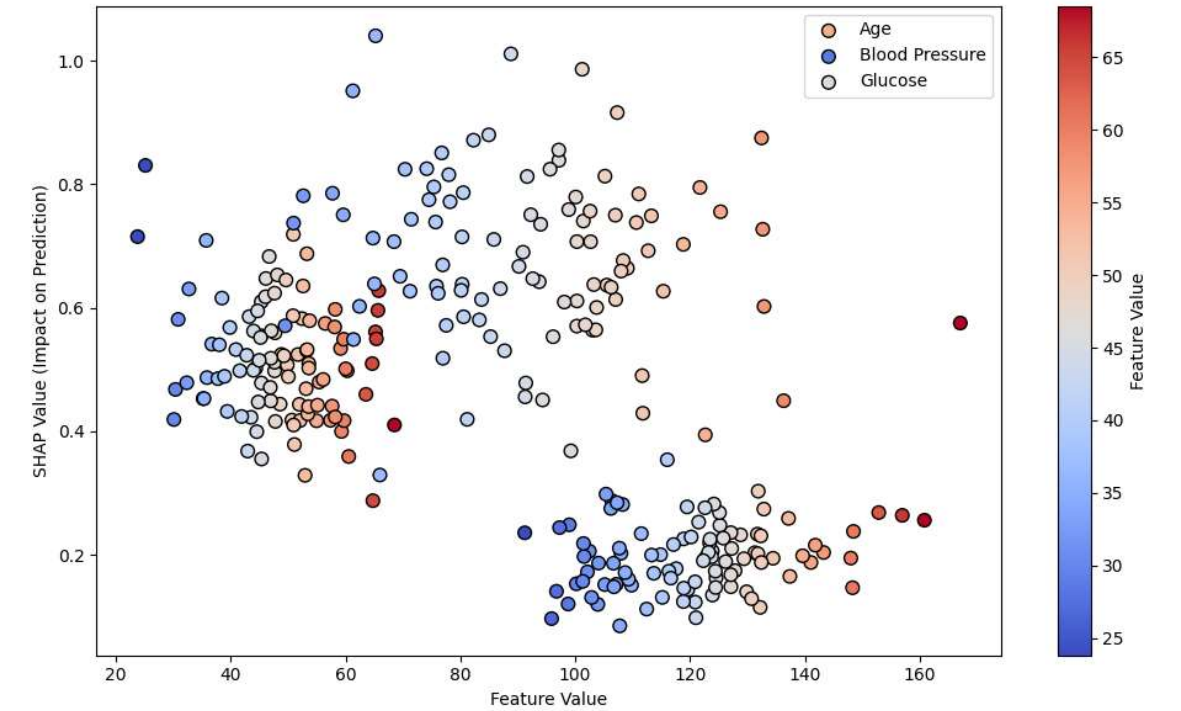


Figure 2. SHAP Values for Risk Prediction in Chronic Disease Diagnosis. The figure illustrates SHAP values for three key features—age, blood pressure, and glucose levels—in a chronic disease risk prediction model. Each dot represents a patient, with the color indicating the feature’s value (e.g., red for higher values and blue for lower values). The x-axis shows the individual feature values, while the y-axis represents the corresponding SHAP values, which measure the impact of each feature on the model's prediction. Positive SHAP values indicate an increased likelihood of disease diagnosis, while negative values suggest a decreased likelihood. This visualization aids in interpreting how the model weighs different features, providing transparency in clinical decision-making.

4. Applications of XAI in Healthcare Diagnostics

XAI continues to make significant contributions to healthcare diagnostics by increasing the transparency, trust, and utility of AI systems. From improving the interpretability of deep learning models in medical imaging to supporting personalized medicine and chronic disease risk prediction, XAI is crucial in bridging the gap between black-box AI models and clinical decision-making.

4.1. Medical Imaging

One of the key areas where XAI has made a profound impact is in medical imaging. Complex AI models, such as CNNs, have revolutionized diagnostic imaging by detecting anomalies like tumors, fractures, and organ irregularities. However, the "black-box" nature of these models presents challenges in clinical adoption, as clinicians need transparency to trust the system’s predictions.

Grad-CAM (Gradient-weighted Class Activation Mapping) remains a widely used XAI method in medical imaging, allowing practitioners to see the specific regions of an image that influence the AI's predictions. Recent advancements in ensemble models combining transfer learning with vision transformers have further enhanced the diagnostic accuracy and explainability of models, especially

in Alzheimer's disease diagnosis. In this approach, a combination of transfer learning (using models like ResNet50 and DenseNet121) with explainability techniques has led to higher accuracy and interpretability, improving diagnostic confidence [18].

Example Use Case: In Alzheimer's diagnosis, Grad-CAM is used to highlight brain regions in MRI scans that are most significant for the model's prediction. This transparency allows neurologists to verify whether the AI's focus corresponds to clinically relevant areas, improving trust in AI-driven diagnostic tools.

4.2. Personalized Medicine in Oncology

Explainable AI is also becoming essential in personalized medicine, particularly in oncology. Personalized treatment plans based on a patient's genetic makeup, tumor characteristics, and clinical history are increasingly being developed using AI models. However, the success of such models depends on their ability to provide clear and understandable treatment recommendations to healthcare providers.

A recent 2024 study used XAI-empowered decision trees for predicting personalized breast cancer treatment options, including hormonal therapies, chemotherapy, and anti-HER2 treatments. The model achieved 99.87% accuracy and improved transparency by explaining which clinical and genetic factors influenced the treatment recommendation. This approach is key to building trust between clinicians and AI systems, as it allows for an understandable rationale behind treatment decisions [19].

Example Use Case: For breast cancer patients, an AI model using clinical and genomic data can recommend specific therapies (e.g., anti-HER2 therapy) based on a decision tree model. With XAI, clinicians can see which genetic mutations or clinical factors led to this recommendation, ensuring that treatment aligns with the patient's specific profile.

4.3. Chronic Disease Risk Prediction

XAI techniques are widely used in predictive analytics for chronic diseases such as diabetes, heart failure, and hypertension. SHAP is particularly useful for identifying the most critical risk factors in predictive models, providing clear explanations for individual patient outcomes. For instance, in predicting heart failure readmissions, SHAP values help explain which factors—such as previous hospitalizations, blood pressure levels, or medication adherence—are driving the model's predictions.

A comprehensive review of AI-based healthcare techniques from 2024 highlights the application of SHAP in identifying key risk factors in disease comorbidity and predicting future health events. XAI techniques like SHAP have allowed clinicians to better understand and intervene in high-risk cases, offering more tailored preventive care [20].

Example Use Case: In predicting heart failure readmission, SHAP values indicate that past hospitalizations and blood pressure levels are the most significant contributors to the model's prediction. By understanding these factors, clinicians can prioritize preventive measures, such as medication adjustments, to reduce the risk of readmission.

4.4. Natural Language Processing (NLP) in Healthcare

Recent research has also focused on Explainable and Interpretable AI (XIAI) in natural language processing (NLP) for healthcare applications. As large language models (LLMs) such as GPT-3 and BERT are increasingly applied to medical tasks—such as analyzing electronic health records (EHRs) or generating diagnostic reports—the need for explainability has become more critical.

A 2024 study introduced the concept of XIAI to distinguish between explainable and interpretable models in healthcare NLP. The study found that attention mechanisms in NLP models can improve interpretability by highlighting relevant portions of clinical text that the model uses to make predictions. This approach enhances model transparency and can improve adoption in critical healthcare applications [21].

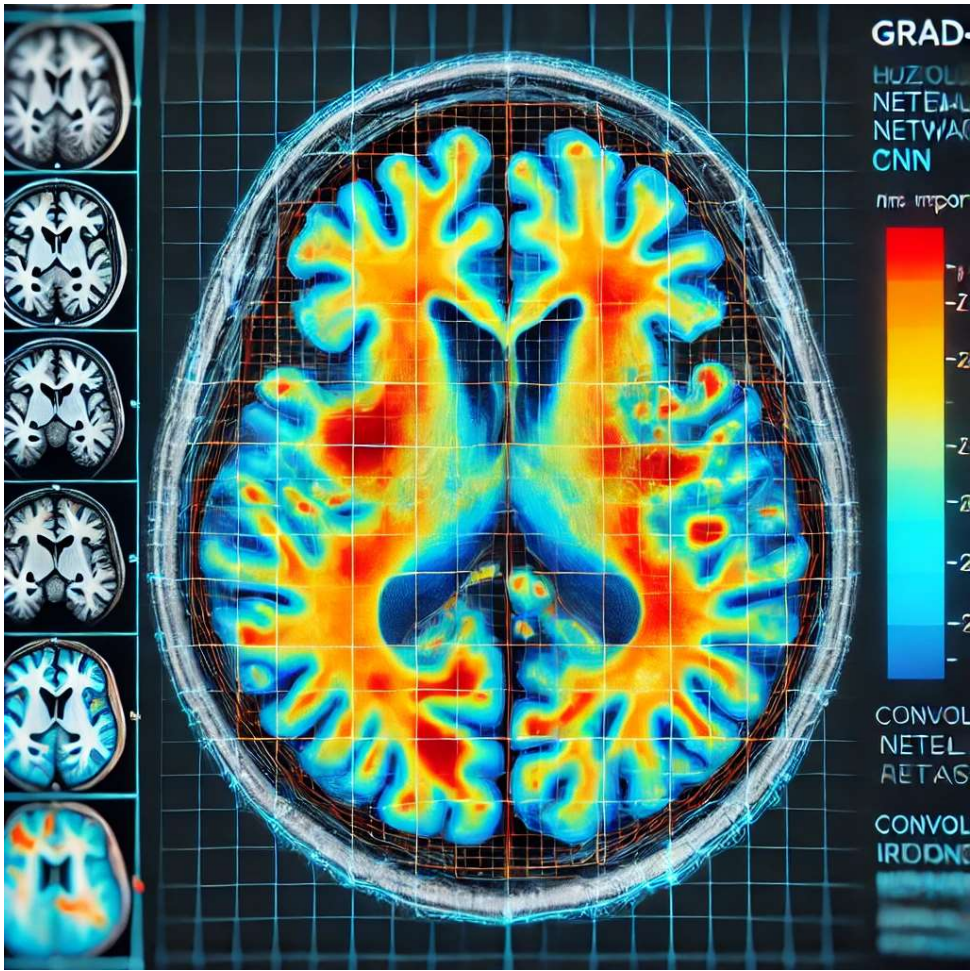
Example Use Case: In analyzing patient EHRs for predicting complications during surgery, an NLP model with attention-based XAI mechanisms can highlight the most relevant parts of the clinical text (e.g., past surgical history or medication use). This allows clinicians to understand how the model arrived at its predictions, increasing confidence in AI-supported decision-making.

Table 3. Applications of XAI Techniques in Healthcare.

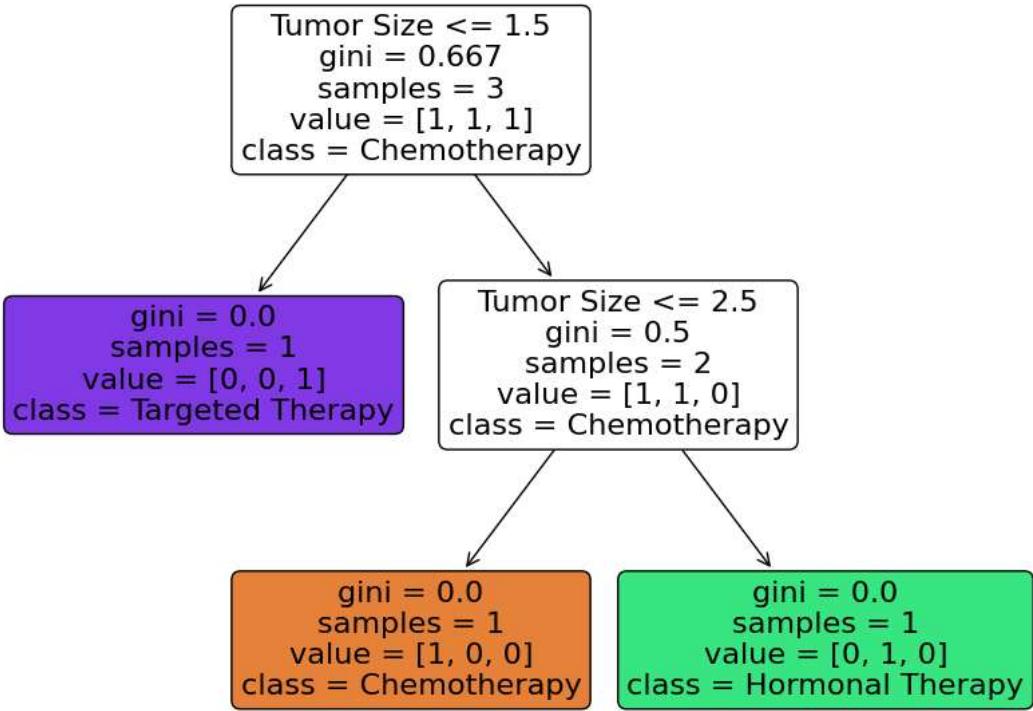
XAI Technique	Application	Purpose	Example Use Case	References
Grad-CAM	Medical Imaging (e.g., Alzheimer's, MRI)	Visualizes regions of medical images important for diagnosis.	Highlighting brain regions in MRI scans for Alzheimer's diagnosis using ensemble models.	[14, 18]
Decision Trees with XAI	Oncology (e.g., Breast Cancer)	Provides transparent, personalized treatment recommendations.	Recommending personalized cancer therapies based on patient genetic markers and clinical data.	[19]
SHAP	Chronic Disease Risk Prediction	Identifies risk factors in predictive models.	Explaining factors like blood pressure and hospitalization history in heart failure readmission prediction.	[20]
Ensemble Transfer Learning & Vision Transformer	Disease Diagnosis	Enhances the accuracy and interpretability of complex models by combining multiple AI techniques.	Alzheimer's disease diagnosis using hybrid models with higher accuracy and clearer interpretability.	[18]
XIAI (Explainable & Interpretable AI)	NLP in Healthcare	Improves model transparency in healthcare NLP tasks for decision-making.	Enhancing interpretability of large language models in personalized medicine and medical task applications.	[21]

The heatmap in **Figure 3A** highlights areas of importance in the brain MRI, where red/orange regions indicate higher relevance in the CNN's decision-making process for diagnosing Alzheimer's. This helps clinicians confirm whether the AI model is focusing on clinically significant regions, improving the transparency of AI predictions. **Figure 3B** illustrates a Decision Tree for breast cancer treatment recommendation, with nodes representing key features such as tumor size, hormone receptor status, HER2 status, and patient age. Each branch leads to a treatment recommendation (chemotherapy, hormonal therapy, or targeted therapy), showing how the AI model arrives at personalized decisions based on clinical and genetic factors. This decision tree is a visual representation of how explainable AI models in oncology provide clear, interpretable pathways for treatment recommendations, making them transparent for clinicians. **Figure 3C** illustrates SHAP values for heart failure readmission prediction, showing the impact of three critical features: age, blood pressure, and previous hospitalizations. Each dot represents a patient, with the position on the x-axis indicating the feature values and the y-axis showing the SHAP values, which measure the contribution of each feature to the model's prediction. Positive SHAP values indicate a higher likelihood of readmission, while negative values suggest a lower likelihood. This visualization helps

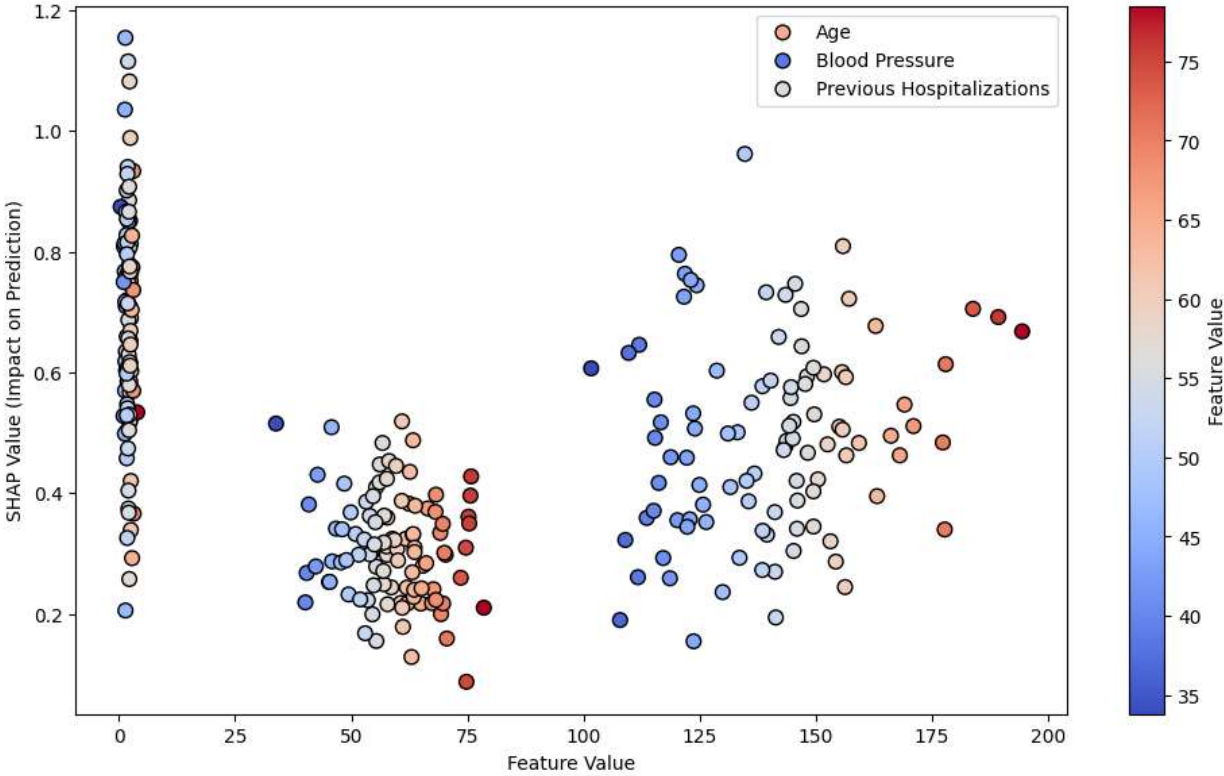
clinicians focus on key risk factors and offers transparency in AI predictions for better decision-making. **Figure 3D** illustrates the Attention Mechanism in NLP for EHR analysis, where the heatmap highlights relevant words in a patient's electronic health record (EHR). This mechanism helps the AI model focus on key terms such as "diabetes" and "blood pressure," which are crucial for making predictions about patient health. The attention scores provide transparency, allowing clinicians to understand which parts of the EHR the model considered most important.



(A)



(B)



(C)



(D)

Figure 3. Illustrative Visualizations of XAI Techniques in Healthcare Diagnostics. (A) Grad-CAM Visualization for Alzheimer's Diagnosis; (B) Decision Tree for Breast Cancer Treatment Recommendation; (C) SHAP Values for Heart Failure Readmission Prediction; (D) Attention Mechanism Highlighting Relevant Words in EHR Analysis.

5. Benefits and Limitations of XAI in Healthcare

XAI has emerged as a pivotal component in the development of AI systems for healthcare diagnostics, offering numerous advantages that improve trust, accountability, and clinical outcomes. However, despite its benefits, XAI also faces significant limitations, particularly when applied to complex medical environments. Below is a detailed examination of both the benefits and limitations of XAI in healthcare.

5.1. Benefits of XAI in Healthcare

5.1.1. Enhanced Trust and Transparency

One of the foremost benefits of XAI in healthcare is its ability to enhance trust and transparency in AI systems. Traditional AI models, such as deep neural networks, often function as "black boxes," making it difficult for clinicians to understand how a diagnosis or treatment recommendation is made. XAI techniques, such as Grad-CAM in medical imaging or SHAP in chronic disease prediction, allow clinicians to visualize or quantify how the model arrived at a particular decision. This interpretability builds trust between healthcare professionals and AI systems, ensuring that decisions are understandable and aligned with clinical knowledge.

For example, in an AI system predicting cancer treatment pathways, a decision tree model enhanced by XAI can clearly display the factors contributing to the recommendation of chemotherapy versus hormonal therapy. Clinicians can see the influence of tumor size, genetic markers, and hormone receptor status, leading to more informed and confident clinical decisions [19].

5.1.2. Improved Regulatory Compliance and Ethical AI

In healthcare, regulatory standards, such as those set by the U.S. Food and Drug Administration (FDA) or European Medicines Agency (EMA), require AI systems to provide interpretable decisions to ensure patient safety and accountability. XAI techniques help meet these regulatory demands by providing insights into how and why certain decisions were made, enabling greater oversight and validation. Explainable models can justify the reasoning behind high-risk decisions such as surgical procedures, drug prescriptions, or disease prognoses, reducing liability risks for healthcare providers and aligning AI systems with ethical standards [18].

5.1.3. Personalized and Precise Patient Care

XAI allows for more personalized and precise patient care by offering detailed insights into individual-level predictions. In personalized medicine, AI models can tailor treatment plans based on patient-specific clinical and genetic data. For example, in precision oncology, XAI-driven models can highlight the genetic markers that led to specific treatment recommendations, allowing clinicians

to tailor therapy to the patient's unique genetic profile. This level of personalization is essential for improving treatment outcomes and minimizing adverse reactions [20].

5.1.4. Facilitates Human-AI Collaboration

XAI promotes better collaboration between AI systems and clinicians by making AI outputs more interpretable and actionable. By providing clear rationales for predictions, XAI enables healthcare professionals to integrate AI insights into their own expertise and decision-making processes. For instance, in radiology, Grad-CAM heatmaps help radiologists see which regions of an MRI the AI model has focused on, allowing them to verify the diagnosis or investigate further [14]. This symbiosis between AI and human expertise leads to more robust diagnostic outcomes.

5.2. Limitations of XAI in Healthcare

5.2.1. Complexity of Interpretability

While XAI aims to make AI models more interpretable, the interpretability itself can sometimes be complex and difficult for non-experts to understand. Techniques like SHAP or LIME, while powerful, produce explanations that may not always be intuitively clear to clinicians without deep technical knowledge. For example, while SHAP values can provide a feature-based explanation for a prediction, understanding the underlying mechanics of Shapley values and how they relate to cooperative game theory can be challenging in a fast-paced clinical setting [13]. Therefore, XAI systems must balance complexity with usability to ensure that the explanations provided are genuinely useful in practice.

5.2.2. Scalability Issues in Large Models

One significant limitation of XAI techniques is their scalability, particularly in large and complex AI models like deep learning networks with millions of parameters. Techniques like Grad-CAM and SHAP, although highly effective for small- to medium-sized models, can become computationally expensive and slow when applied to large-scale networks used in medical diagnostics. This can limit the practical deployment of XAI in real-time clinical environments, where speed and accuracy are crucial [20].

5.2.3. Limitations in Explaining Certain AI Models

Not all AI models lend themselves easily to explanation through XAI techniques. For instance, while models like decision trees or linear regression are inherently interpretable, more complex models such as deep reinforcement learning or RNNs used for time-series medical data often remain difficult to explain using current XAI methods. This presents a challenge in certain healthcare applications, such as long-term patient monitoring or predictive modeling based on temporal data, where clinicians may struggle to interpret and validate the model's decisions.

5.2.4. Risk of Over-Simplification

In efforts to make AI models more interpretable, there is a risk that XAI techniques may oversimplify the underlying logic, potentially leading to inaccurate or misleading conclusions. For example, in some cases, post-hoc explanation methods like LIME may provide locally accurate explanations that do not fully reflect the global behavior of the model. This could lead clinicians to make decisions based on incomplete or misleading insights, which could compromise patient care [15].

5.2.5. Ethical and Bias Concerns

XAI does not fully eliminate ethical concerns around AI systems, particularly regarding bias. While XAI can make the decision-making process more transparent, it does not inherently address the issue of biased data or biased decision-making within the models themselves. If the underlying

AI system is trained on biased datasets, even the explanations provided by XAI may reflect and perpetuate these biases, potentially leading to discriminatory outcomes in healthcare, particularly for underrepresented groups [21].

In summary, the integration of XAI in healthcare brings a multitude of benefits, from increasing trust and transparency to supporting personalized patient care and ensuring regulatory compliance. However, the limitations of XAI, including interpretability challenges, scalability issues, and ethical concerns, need to be carefully considered in its adoption. As AI continues to evolve, so too must the techniques for making these systems explainable, ensuring that they are both powerful and practical for real-world clinical applications.

Table 4 summarizes the primary benefits and limitations of applying XAI in healthcare diagnostics. It contrasts how XAI improves trust, transparency, regulatory compliance, and personalized care, while also addressing challenges such as scalability, interpretability complexity, and potential biases.

Table 4. Summary of Benefits and Limitations of XAI in Healthcare.

Aspect	Benefits	Limitations	References
Trust and Transparency	Enhances clinician trust with interpretable decisions.	Interpretability complexity can hinder understanding for non-experts.	[13, 14]
Regulatory Compliance	Supports regulatory requirements for safe and accountable AI.	Not all AI models can be easily explained (e.g., RNNs).	[18]
Personalized Care	Allows more precise and tailored treatment plans for individual patients.	Over-simplification in post-hoc methods may mislead clinicians.	[19, 20]
Human-AI Collaboration	Facilitates better collaboration between AI systems and clinicians.	Scalability issues in larger models affect real-time performance.	[21]
Ethical Considerations	Helps address bias and fairness by providing insights into model decisions.	XAI may still perpetuate biases present in the training data.	[15, 21]

Key Points:

- This table clearly outlines both the advantages and challenges associated with using XAI in healthcare, making it easier for readers to compare both sides.
- The benefits cover trust-building, regulatory support, and personalized care, while the limitations focus on interpretability, scalability, and ethical concerns.

Figure 4A workflow diagram illustrates how XAI integrates into healthcare diagnostics. It begins with patient data input (e.g., medical images or EHR), followed by AI model processing, the application of XAI techniques (e.g., Grad-CAM, SHAP), the generation of explainable outputs (heatmaps, feature importance scores), and concludes with human decision-making, where clinicians review the explainable outputs to make final diagnoses or treatment decisions.

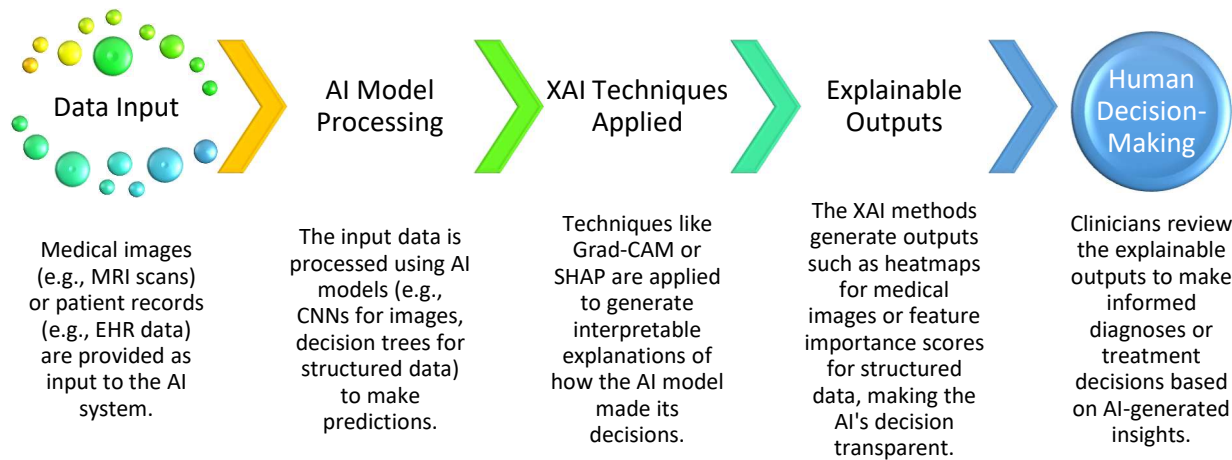
Key Points:

- The workflow shows the step-by-step process in which XAI techniques are embedded in the healthcare decision-making pipeline.
- It emphasizes how XAI provides interpretable insights that allow clinicians to verify AI-generated recommendations, improving trust and reliability in AI-based diagnostics.

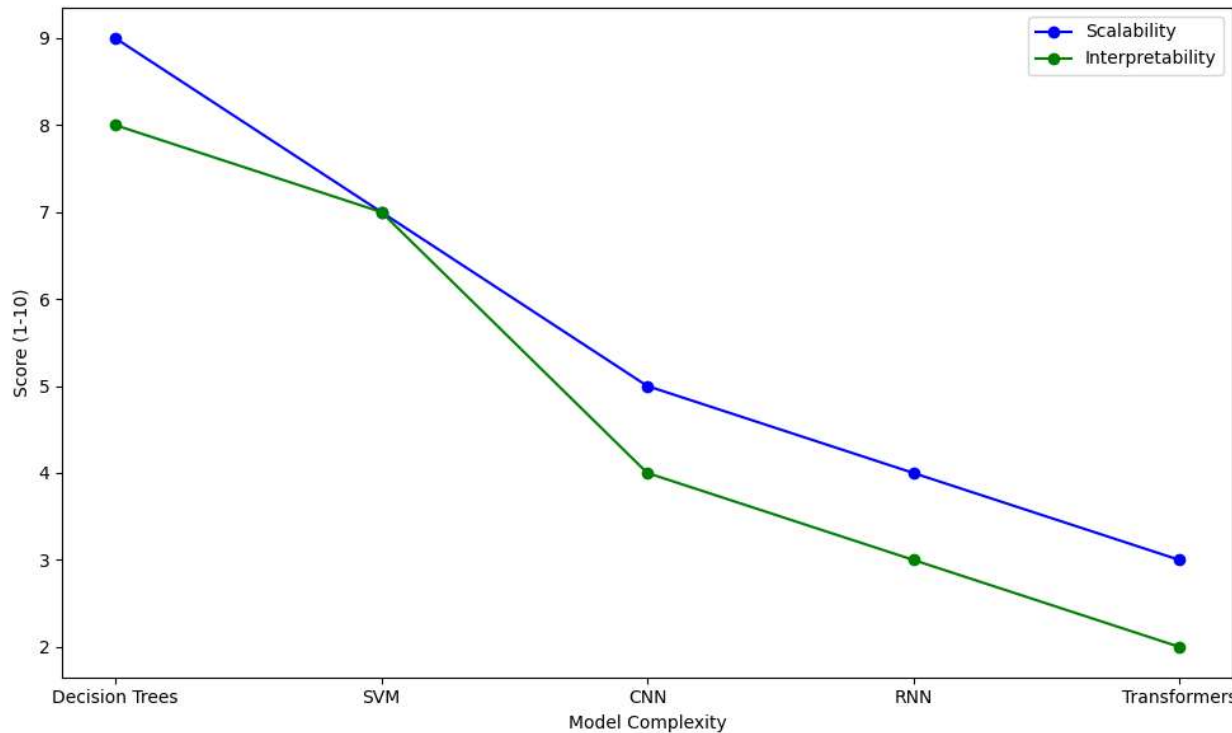
Figure 4B graph shows the relationship between model complexity and the scalability and interpretability of XAI techniques. As AI models become more complex (e.g., transitioning from decision trees to CNNs or transformers), the scalability and interpretability of XAI techniques decrease. The graph highlights the performance trade-offs between maintaining explainability and handling larger, more complex models.

Key Points:

- Simpler models (e.g., decision trees) are easier to interpret but may not scale well for large datasets or complex tasks.
- More complex models (e.g., CNNs, Transformers) offer higher accuracy but are harder to interpret and require more computational resources, limiting their scalability.
- The graph helps visualize how XAI methods like SHAP and Grad-CAM perform in relation to model complexity and provides insight into the trade-offs between interpretability and computational efficiency.



(A)



(B)

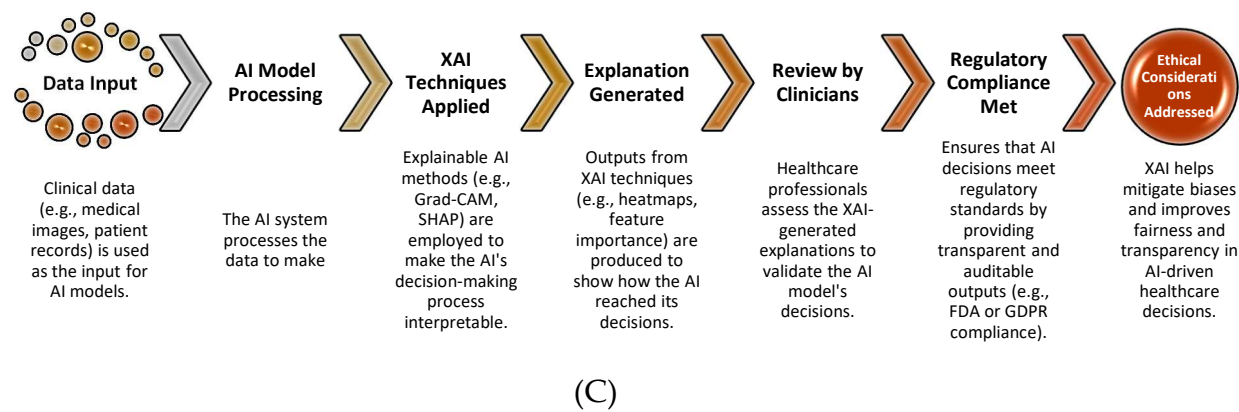


Figure 4. Illustrative Visuals for Benefits and Limitations of XAI in Healthcare. (A) Workflow of XAI in Healthcare Diagnostics; (B) Scalability of XAI Techniques vs Model Complexity; (C) Flowchart XAI Impact on Regulatory Compliance and Ethical Considerations.

Figure 4C flowchart illustrates the role of XAI in ensuring regulatory compliance and addressing ethical concerns such as bias and transparency in healthcare. The diagram shows how data is processed through AI models, explained using XAI techniques, and reviewed by clinicians, ultimately achieving regulatory compliance (e.g., FDA, GDPR) and addressing ethical considerations (e.g., bias detection).

Key Points:

- XAI techniques play a vital role in aligning AI models with regulatory frameworks, ensuring transparency and auditability.
- The flowchart highlights how ethical concerns, such as fairness and bias in AI decisions, can be mitigated by applying XAI methods, ultimately leading to safer and more accountable healthcare systems.

General Observations for All Visual Elements:

- **Context:** These visuals collectively explain how XAI contributes to improving healthcare diagnostics while simultaneously addressing its limitations, including regulatory and ethical challenges.
- **Application:** They highlight how XAI fits into healthcare workflows, balancing the benefits of interpretability with the practical challenges of scalability and regulatory requirements.
- **Importance:** The visuals emphasize how XAI techniques such as Grad-CAM, SHAP, and LIME help bridge the gap between AI model complexity and the need for transparent, reliable decision-making in clinical environments.

6. Ethical and Regulatory Considerations in XAI for Healthcare

As AI continues to be adopted in healthcare, ensuring that AI systems are ethical, transparent, and compliant with regulations has become a critical concern. XAI plays a crucial role in addressing both ethical and regulatory challenges by making the decision-making processes of AI models more understandable. Below is an exploration of the ethical implications and regulatory requirements for XAI in healthcare.

6.1. Ethical Considerations in XAI

6.1.1. Addressing Bias and Fairness

One of the most pressing ethical concerns in healthcare AI is the potential for biased or unfair decision-making. If an AI model is trained on data that underrepresents certain populations or is

skewed toward specific demographic groups, it may produce biased results. For instance, an AI system trained primarily on data from one ethnic group may not perform as well on patients from other ethnicities. XAI techniques, such as SHAP or LIME, provide visibility into which features the AI model considers important for making predictions, helping to identify and correct for bias.

In recent studies, XAI has been used to expose bias in models used for predicting disease risks in minority populations. By making the decision-making process transparent, XAI allows researchers and clinicians to scrutinize whether certain features (e.g., race, gender) disproportionately affect the AI's decisions and adjust the model accordingly [21].

6.1.2. Ensuring Transparency

Transparency is a key ethical concern in healthcare AI. Patients and clinicians must understand how AI systems arrive at their decisions, especially when these decisions have life-or-death consequences, such as diagnosing cancer or recommending surgical interventions. XAI techniques provide the necessary tools to increase transparency by explaining the reasoning behind AI-driven decisions. In turn, this fosters trust between clinicians, patients, and AI systems.

Grad-CAM, for example, can show exactly which parts of a medical image a model focuses on when diagnosing a condition like Alzheimer's. This helps radiologists verify whether the AI's focus aligns with known clinical indicators, improving the reliability and ethical integrity of AI-assisted diagnostics [14].

6.1.3. Accountability and Trust

In healthcare, where incorrect diagnoses or treatment recommendations can lead to severe consequences, accountability is paramount. XAI enhances accountability by allowing clinicians to understand and question the AI's decision-making process, ensuring that decisions are grounded in clinical reasoning. This is particularly important for healthcare providers, as they are ultimately responsible for patient outcomes. If an AI system's decisions are opaque, clinicians may hesitate to trust it, which can hinder adoption. XAI helps bridge this gap by making AI systems more trustworthy.

A key application is in personalized medicine, where clinicians rely on AI to make recommendations for tailored treatments. By using XAI methods like decision trees or SHAP, clinicians can understand which factors influenced the treatment suggestion, thereby making the AI system more accountable and improving its integration into healthcare workflows [19].

6.2. Regulatory Considerations in XAI

6.2.1. Meeting Regulatory Standards

AI systems in healthcare must comply with various regulatory standards to ensure safety, accuracy, and transparency. Regulatory bodies such as the U.S. Food and Drug Administration (FDA) and the European Medicines Agency (EMA) require AI models to be transparent and auditable. XAI techniques are crucial for meeting these standards by providing interpretable outputs that can be scrutinized by regulators.

For example, the FDA's AI/ML-based Software as a Medical Device (SaMD) Action Plan emphasizes the need for transparency and continuous learning in AI systems. XAI methods can provide the necessary insights into how an AI system makes decisions, allowing regulators to assess the model's safety and reliability before approving it for clinical use [18].

6.2.2. Audibility and Validation

XAI not only improves transparency but also enhances the audibility of AI systems in healthcare. Regulators require that AI systems be auditable, meaning that their decisions can be tracked and validated. This is especially important in sensitive areas like diagnostics or treatment

recommendations, where an audit trail ensures that decisions are well-founded and based on sound medical principles.

By using XAI, healthcare providers can maintain an audit trail of the AI’s decision-making process. For instance, if an AI model recommends a particular treatment, XAI can highlight the key factors that influenced the decision. This audit trail can be reviewed by both regulators and clinicians to ensure the AI system is making decisions aligned with regulatory standards.

6.2.3. Data Privacy and Security

Data privacy is a key regulatory concern, particularly with the advent of laws like General Data Protection Regulation (GDPR) in Europe and the Health Insurance Portability and Accountability Act (HIPAA) in the United States. These regulations require that patient data be handled securely and that any decisions made using AI systems are explainable, especially when personal data is involved. XAI techniques ensure that AI models comply with these regulations by providing clear explanations about how sensitive data is used in decision-making.

XAI helps balance the need for data-driven insights with the requirement to protect patient privacy. For example, if a model uses patient demographics to make a prediction, XAI can ensure that this data is used appropriately and that the model’s decision-making process complies with relevant privacy laws [20].

6.3. The Future of XAI in Ethical and Regulatory Frameworks

As AI continues to evolve, so too must the ethical and regulatory frameworks surrounding it. In the future, regulatory bodies may require not only explainable AI models but also continuous monitoring of AI systems to ensure they remain transparent and unbiased over time. As AI systems learn and evolve with new data, it will be critical to maintain their explainability to ensure compliance with ethical and regulatory standards.

Future frameworks may also prioritize responsible AI principles, which emphasize fairness, transparency, and accountability. XAI will play an integral role in ensuring that AI systems in healthcare align with these principles, ensuring that they contribute to better patient outcomes while safeguarding ethical integrity.

In summary, XAI serves as a vital tool in addressing the ethical and regulatory challenges posed by AI in healthcare. By improving transparency, mitigating bias, and providing accountability, XAI helps build trust in AI systems while ensuring they comply with regulatory requirements. As AI continues to integrate into healthcare, XAI will be key in shaping ethical AI practices and ensuring that AI systems are both effective and responsible.

Table 5 summarizes the primary ethical concerns in applying AI in healthcare, such as bias, transparency, accountability, and data privacy. It highlights how specific XAI techniques (like SHAP and Grad-CAM) address these concerns. For example, SHAP helps identify biased feature attribution in patient outcomes, while Grad-CAM enhances transparency in medical image analysis by providing visual explanations of AI decisions. This table serves as a quick reference for understanding how XAI mitigates ethical risks and improves trust in AI-driven decisions.

Table 5. Ethical Considerations in XAI for Healthcare.

Ethical Aspect	Description	XAI Technique Addressing the Concern	References
Bias and Fairness	Identifies biased or unfair decision-making in AI models due to imbalanced data or flawed feature selection.	SHAP for feature attribution, LIME for local explanations.	[21]
Transparency	Ensures that clinicians and patients can understand how AI systems arrive at their decisions.	Grad-CAM for visual explanations in medical imaging.	[14]

Accountability	Improves trust by providing clinicians with the ability to trace and verify AI decisions.	Decision Trees, SHAP for global interpretability. [19]
Data Privacy	Ensures that patient data is handled securely and used appropriately in AI models.	Data usage transparency through XAI outputs (SHAP, LIME). [20]

Key Points:

- Ethical challenges are directly linked to the decision-making process of AI models.
- XAI techniques provide transparency, fairness, and auditability in healthcare applications.

Figure 5A flowchart outlines the critical steps in ensuring that XAI models meet healthcare regulatory standards. The workflow begins with input data (e.g., medical images or patient records), processed through AI models (CNNs or decision trees). XAI techniques (e.g., Grad-CAM, SHAP) are applied to generate explainable outputs, such as heatmaps or feature importance scores, which are reviewed by clinicians. The final stage includes validation and auditing by regulators, ensuring compliance with frameworks such as FDA, GDPR, and HIPAA.

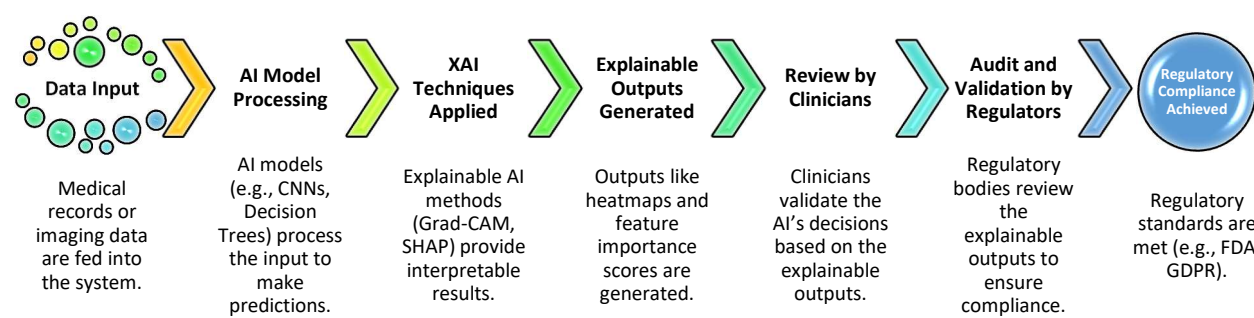
Key Points:

- XAI is instrumental in achieving transparency and explainability for regulatory audits.
- Clinicians and regulators can scrutinize the model’s decisions, ensuring that they are aligned with clinical and legal standards.

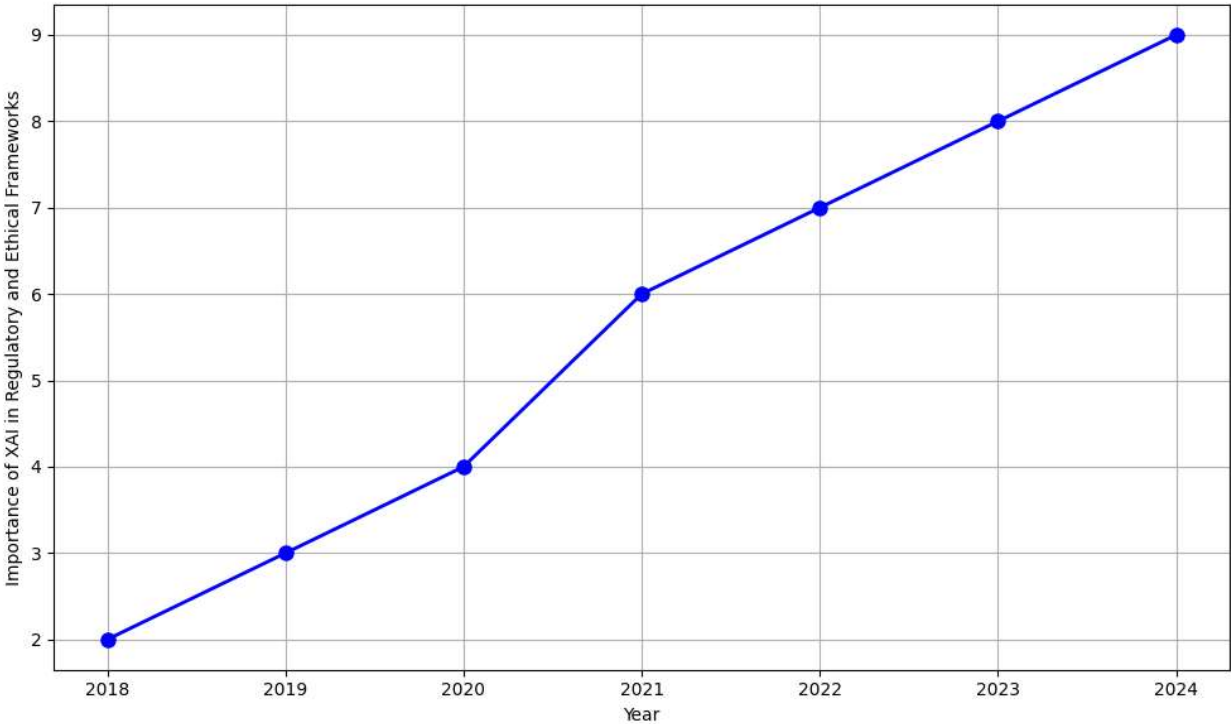
Figure 5B line graph shows the growing importance of XAI in healthcare from 2018 to 2024. The upward trend reflects the increasing demand for explainability as AI systems become more prevalent in medical practice. Regulatory bodies such as the FDA and EMA are progressively requiring AI systems to offer interpretability to meet safety, fairness, and transparency criteria. The graph highlights that, over time, XAI is playing a more significant role in aligning AI technologies with ethical and regulatory standards.

Key Points:

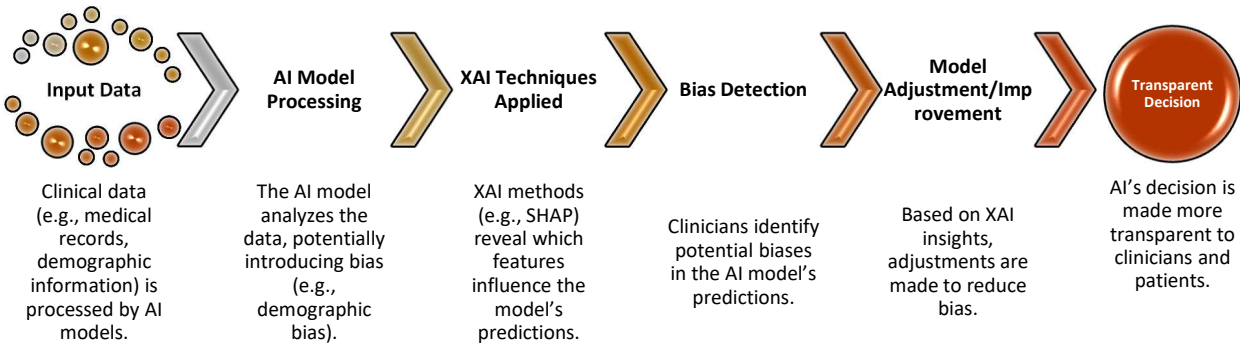
- The rising curve represents the heightened focus on XAI by regulators and healthcare institutions.
- XAI is critical in achieving regulatory approval for AI-based healthcare applications as explainability becomes a primary concern.



(A)



(B)



(C)

Figure 5. Illustrative Visuals for Ethical and Regulatory Considerations in XAI for Healthcare. (A) Workflow of XAI's Role in Regulatory Compliance; (B) Impact of XAI on Ethical and Regulatory Frameworks (2018-2024); (C) XAI for Bias Mitigation and Transparency in Decision-Making.

Figure 5C diagram illustrates how XAI techniques, like SHAP, help mitigate bias and improve transparency in healthcare AI systems. The workflow shows how data, when processed by an AI model, may introduce bias. XAI techniques are then applied to uncover the factors that influence the AI's decision, helping clinicians detect and address any biases. By adjusting the model based on XAI feedback, clinicians can ensure fair and transparent decisions for patient care.

Key Points:

- XAI helps expose biases that may arise in AI-driven predictions (e.g., racial or gender bias).
- The diagram showcases how XAI can be integrated into AI workflows to ensure that decision-making is fair and transparent.

General Observations for All Visual Elements:

- **Context:** These visual elements demonstrate the dual role of XAI in ensuring compliance with both ethical standards (like fairness, transparency, and accountability) and regulatory requirements (such as those imposed by the FDA, GDPR, and HIPAA).

- **Application:** They highlight that XAI techniques are essential in providing explainability, which is crucial for the widespread adoption of AI systems in clinical environments.
- **Importance:** The visuals emphasize the growing reliance on XAI to ensure that AI models in healthcare are interpretable, compliant, and ethically aligned with patient safety and care.

7. Discussion

XAI has been a game-changer in healthcare diagnostics, offering insights into the often opaque decision-making processes of advanced AI models. While XAI techniques like SHAP, Grad-CAM, and LIME have bridged the gap between AI model outputs and clinician understanding, the implementation of XAI in real-world healthcare systems still faces significant challenges. This discussion addresses the core benefits and limitations of XAI, focusing on trust, accountability, bias mitigation, scalability, and regulatory compliance.

7.1. Balancing Accuracy and Interpretability

One of the most persistent challenges in applying XAI in healthcare is the trade-off between the accuracy of complex models (like deep learning models) and their interpretability. Advanced AI models, such as CNNs and transformers, are known for their high accuracy in tasks like medical imaging analysis, but they often function as "black boxes," making their decisions hard to explain and understand [14]. On the other hand, simpler models like decision trees are more interpretable but may not offer the same level of accuracy for complex tasks [22]. XAI techniques such as SHAP have been introduced to provide local explanations for individual predictions, enabling clinicians to see which features most influenced the model's decision [23].

Despite these advancements, current XAI methods are computationally expensive, which limits their application in real-time diagnostics where speed is critical [24]. Moreover, scalability remains a challenge, as models increase in complexity, and ensuring that XAI can function efficiently in time-sensitive clinical settings requires further research [25].

7.2. Trust and Accountability in AI-Driven Healthcare

Building trust between clinicians and AI systems is crucial for the widespread adoption of AI in healthcare. Studies show that clinicians are more likely to use AI systems if they can understand and verify the reasoning behind the decisions. Tools like Grad-CAM have been particularly helpful in medical imaging, allowing clinicians to visualize which parts of an MRI or CT scan the model focused on for diagnosis [14]. This has improved the transparency of AI in radiology and increased clinician confidence in AI-based diagnostic tools [26].

However, interpretability alone does not ensure accountability. AI systems are not immune to errors, and when these errors occur, particularly in critical areas like diagnosis or treatment recommendations, clinicians need to trace the logic behind the decision to determine whether it was valid [15]. XAI techniques allow this, but further work is needed to define clear accountability frameworks, particularly in cases where AI-driven decisions lead to adverse outcomes [2].

7.3. Ethical Considerations and Bias Mitigation

AI models trained on biased datasets can perpetuate and even exacerbate existing health disparities, leading to unequal treatment of patients. This is especially concerning in healthcare, where underrepresented populations may be at risk of receiving suboptimal care due to biased AI predictions [27]. XAI techniques offer a way to identify and mitigate bias by providing transparency in model predictions. For instance, SHAP can highlight whether certain demographic features (such as race or gender) disproportionately influence the model's decision-making process [20]. This allows for adjustments to be made to the model or the data to ensure fairer outcomes.

However, XAI cannot eliminate all biases inherent in data. The presence of biases in medical datasets remains a significant issue that XAI alone cannot solve. Ethical guidelines and rigorous data collection protocols are essential to ensure that AI systems do not reinforce discriminatory practices

[28]. Future research should focus on the development of ethical AI frameworks that prioritize fairness and accountability while still leveraging the advantages of XAI techniques [29].

7.4. Regulatory Compliance and Legal Implications

The increasing use of AI in healthcare has attracted attention from regulatory bodies, such as the FDA and the European Medicines Agency (EMA), which emphasize the need for transparency and explainability in AI systems [18]. XAI plays a critical role in ensuring that AI models meet these regulatory requirements by providing interpretable outputs that can be scrutinized for safety and effectiveness [30]. Legal frameworks are also beginning to address the role of AI in clinical decision-making, particularly in defining who is accountable when an AI system contributes to a medical error [31].

Additionally, compliance with privacy regulations, such as GDPR and HIPAA, is crucial when using patient data in AI systems [32]. XAI systems must not only explain their decisions but also ensure that sensitive patient data is protected and used in compliance with legal standards. As regulatory frameworks evolve, XAI will need to adapt to provide explanations that meet these rigorous standards while maintaining patient confidentiality [25].

7.5. The Future of XAI in Healthcare

The future of XAI in healthcare lies in the development of hybrid models that combine the interpretability of simpler models with the accuracy of more complex deep learning systems [22]. Explainable-by-design AI systems, which inherently incorporate transparency, may reduce the reliance on post-hoc explanation techniques like SHAP and Grad-CAM [33]. These models are more likely to be integrated into real-time clinical workflows, where speed and accuracy are paramount [34].

Moreover, there is a growing need for intuitive user interfaces that allow clinicians to interact with XAI models easily, facilitating the integration of AI into day-to-day clinical decision-making [2]. As the field of AI continues to evolve, the development of scalable, efficient XAI methods will be key to ensuring that AI is trusted, reliable, and ethically sound in healthcare.

In summary, this discussion highlights the critical role XAI plays in healthcare, improving transparency, trust, and accountability while addressing bias and regulatory challenges. However, much work remains to be done, particularly in balancing accuracy with interpretability, mitigating bias, and developing scalable solutions for real-time clinical use. As AI continues to transform healthcare, XAI will be indispensable in ensuring that these technologies are not only powerful but also ethically aligned with patient needs.

8. Conclusion

XAI is transforming healthcare diagnostics by addressing one of the most pressing challenges in artificial intelligence—understanding and interpreting the complex decision-making processes of AI models. By making AI models more transparent and interpretable, XAI enhances trust, accountability, and collaboration between clinicians and AI systems. However, despite these promising advancements, significant challenges remain, particularly around the trade-off between accuracy and interpretability, scalability of XAI techniques, bias mitigation, and regulatory compliance.

One of the key benefits of XAI is its ability to improve clinician trust in AI systems by providing interpretable outputs that explain how and why a model arrived at a particular diagnosis or recommendation. Techniques such as Grad-CAM, SHAP, and LIME have enabled clinicians to visualize and understand AI decisions, fostering greater confidence in AI-assisted diagnostics. This is particularly important in fields such as radiology and oncology, where AI systems are increasingly relied upon to detect diseases early and recommend personalized treatments. Nevertheless, the complexity of interpretability, especially in deep learning models like CNNs and transformers, remains a critical challenge.

Another advantage of XAI is its potential to enhance regulatory compliance. As AI becomes more integral to clinical workflows, regulatory bodies such as the FDA and EMA are emphasizing the importance of explainability in AI models. XAI techniques help ensure that AI systems can meet these regulatory requirements, providing transparency and auditability that are essential for safe and ethical deployment in healthcare. However, scalability continues to pose a problem, as more complex models often require extensive computational resources to generate interpretable explanations, limiting their real-time application in clinical settings.

Ethical considerations, particularly around bias and fairness, also play a central role in the XAI discussion. While XAI techniques provide a pathway for identifying and mitigating bias, they cannot fully eliminate biases embedded in datasets. As AI models are only as good as the data they are trained on, biased datasets will inevitably lead to biased predictions. This highlights the need for more rigorous data collection processes and continuous bias audits throughout the AI lifecycle. Regulatory frameworks will need to evolve to address these issues, ensuring that AI systems are not only effective but also fair and equitable for all patient groups.

Looking to the future, the development of hybrid models that combine the high accuracy of deep learning with the interpretability of simpler models may offer a promising solution to some of these challenges. Moreover, advancements in explainable-by-design AI models, which are inherently interpretable without requiring post-hoc explanation techniques, could further enhance the applicability of AI in healthcare. As AI technology continues to evolve, XAI will remain an essential component, ensuring that AI systems are transparent, accountable, and aligned with the ethical standards required in healthcare.

In conclusion, while XAI has made significant strides in making AI-driven healthcare diagnostics more interpretable and trustworthy, it is clear that ongoing research is needed to address the remaining challenges. Future efforts should focus on improving the scalability of XAI techniques, reducing computational costs, and enhancing bias detection. At the same time, collaboration between AI developers, clinicians, and regulators will be crucial to ensuring that XAI systems are seamlessly integrated into healthcare, ultimately improving patient outcomes and ensuring the ethical deployment of AI technologies.

Funding: This research received no external funding.

Compliance with Ethical Standards: This article does not involve any studies conducted by the authors that included human participants.

Acknowledgments: The completion of this research work was made possible through the collaborative efforts and dedication of a multidisciplinary team. We extend our sincere appreciation to each member for their invaluable contributions.

Conflicts of Interest: The authors declare no conflicts of interest.

References

1. Litjens, G., et al., *A survey on deep learning in medical image analysis*. Medical image analysis, 2017. **42**: p. 60-88.
2. Topol, E.J., High-performance medicine: the convergence of human and artificial intelligence. *Nature medicine*, 2019. **25**(1): p. 44-56.
3. Zech, J.R., et al., Variable generalization performance of a deep learning model to detect pneumonia in chest radiographs: a cross-sectional study. *PLoS medicine*, 2018. **15**(11): p. e1002683.
4. Samek, W., et al., *Explainable AI: interpreting, explaining and visualizing deep learning*. Vol. 11700. 2019: Springer Nature.
5. Esteva, A., et al., Dermatologist-level classification of skin cancer with deep neural networks. *nature*, 2017. **542**(7639): p. 115-118.
6. Choi, E., et al. Doctor ai: Predicting clinical events via recurrent neural networks. in *Machine learning for healthcare conference*. 2016. PMLR.
7. Cruz, J.A. and D.S. Wishart, *Applications of machine learning in cancer prediction and prognosis*. *Cancer informatics*, 2006. **2**: p. 117693510600200030.

8. Quinlan, J.R., *Induction of decision trees*. Machine learning, 1986. 1: p. 81-106.
9. Ghassemi, M., et al., *A review of challenges and opportunities in machine learning for health*. AMIA Summits on Translational Science Proceedings, 2020. 2020: p. 191.
10. Wang, F., L.P. Casalino, and D. Khullar, *Deep learning in medicine—promise, progress, and challenges*. JAMA internal medicine, 2019. 179(3): p. 293-294.
11. Obermeyer, Z., et al., Dissecting racial bias in an algorithm used to manage the health of populations. Science, 2019. 366(6464): p. 447-453.
12. Selvaraju, R.R., et al. Grad-cam: Visual explanations from deep networks via gradient-based localization. in Proceedings of the IEEE international conference on computer vision. 2017.
13. Lundberg, S.M., et al., From local explanations to global understanding with explainable AI for trees. Nature machine intelligence, 2020. 2(1): p. 56-67.
14. Selvaraju, R.R., et al., *Grad-CAM: visual explanations from deep networks via gradient-based localization*. International journal of computer vision, 2020. 128: p. 336-359.
15. Ribeiro, M.T., S. Singh, and C. Guestrin. "Why should i trust you?" Explaining the predictions of any classifier. in Proceedings of the 22nd ACM SIGKDD international conference on knowledge discovery and data mining. 2016.
16. Shrikumar, A., P. Greenside, and A. Kundaje. Learning important features through propagating activation differences. in International conference on machine learning. 2017. PMIR.
17. Quinlan, J.R., C4. 5: programs for machine learning. 2014: Elsevier.
18. Poonia, R.C. and H.A. Al-Shaikh, Ensemble approach of transfer learning and vision transformer leveraging explainable AI for disease diagnosis: An advancement towards smart healthcare 5.0. Computers in Biology and Medicine, 2024. 179: p. 108874.
19. Reena R Lokare, J.W., Sunita Patil, Ganesh Wadmare, Darshan Patil, *Transparent precision: Explainable AI empowered breast cancer recommendations for personalized treatment*. IAES International Journal of Artificial Intelligence, 2024. 13(3): p. 2694-2694.
20. Kalra, N., P. Verma, and S. Verma, *Advancements in AI based healthcare techniques with FOCUS ON diagnostic techniques*. Computers in Biology and Medicine, 2024. 179: p. 108917.
21. Huang, G., et al., From explainable to interpretable deep learning for natural language processing in healthcare: How far from reality? Computational and Structural Biotechnology Journal, 2024.
22. Loh, H.W., et al., Application of explainable artificial intelligence for healthcare: A systematic review of the last decade (2011–2022). Computer Methods and Programs in Biomedicine, 2022. 226: p. 107161.
23. Lundberg, S., *A unified approach to interpreting model predictions*. arXiv preprint arXiv:1705.07874, 2017.
24. Doshi-Velez, F. and B. Kim, *Towards a rigorous science of interpretable machine learning*. arXiv preprint arXiv:1702.08608, 2017.
25. Samek, W., Explainable artificial intelligence: Understanding, visualizing and interpreting deep learning models. arXiv preprint arXiv:1708.08296, 2017.
26. Keller, J., et al., Exercise leads to metabolic changes associated with improved strength and fatigue in people with MS. Annals of Clinical and Translational Neurology, 2021. 8(6): p. 1308-1317.
27. Mehrabi, N., et al., *A survey on bias and fairness in machine learning*. ACM computing surveys (CSUR), 2021. 54(6): p. 1-35.
28. Gianfrancesco, M.A., et al., Potential biases in machine learning algorithms using electronic health record data. JAMA internal medicine, 2018. 178(11): p. 1544-1547.
29. Rajkomar, A., J. Dean, and I. Kohane, *Machine learning in medicine*. New England Journal of Medicine, 2019. 380(14): p. 1347-1358.
30. Watson, D.S., et al., Clinical applications of machine learning algorithms: beyond the black box. Bmj, 2019. 364.
31. He, J., et al., The practical implementation of artificial intelligence technologies in medicine. Nature medicine, 2019. 25(1): p. 30-36.
32. Al-Rubaie, M. and J.M. Chang, *Privacy-preserving machine learning: Threats and solutions*. IEEE Security & Privacy, 2019. 17(2): p. 49-58.
33. Christoph, M., *Interpretable machine learning: A guide for making black box models explainable*. 2020: Leanpub.
34. Parikh, R.B., S. Teeple, and A.S. Navathe, *Addressing bias in artificial intelligence in health care*. Jama, 2019. 322(24): p. 2377-2378.

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