

Interpretable Models for Healthcare: A Comparative Analysis of Explainable Machine Learning Approaches

Anand R. Mehta

ABSTRACT

As machine learning models become increasingly prevalent in healthcare settings, the need for interpretability and transparency in these models is paramount to gain trust from healthcare practitioners, ensure patient safety, and facilitate effective decision-making. This study presents a comprehensive comparative analysis of various explainable machine learning (XAI) approaches applied to healthcare datasets. The objective is to evaluate and compare the interpretability, accuracy, and utility of different XAI techniques to aid in the selection of suitable models for healthcare applications. The research employs a diverse set of healthcare datasets, encompassing various medical domains such as diagnostic imaging, electronic health records, and patient outcomes. We investigate popular XAI techniques, including but not limited to LIME (Local Interpretable Model-agnostic Explanations), SHAP (Shapley Additive exPlanations), decision trees, and rule-based models. Each method is assessed based on its ability to provide meaningful explanations for model predictions, its accuracy in capturing complex medical relationships, and its utility in aiding healthcare professionals in understanding and trusting the model outputs.

Our findings reveal that different XAI approaches exhibit varying levels of interpretability and performance across healthcare datasets. While certain methods excel in providing local explanations for specific predictions, others demonstrate superior global interpretability. Additionally, the study explores the trade-offs between interpretability and model accuracy, shedding light on the balance required for practical implementation in healthcare scenarios. Furthermore, we discuss the ethical implications of deploying opaque models in healthcare and emphasize the importance of considering the societal impact of algorithmic decision-making. The study concludes by providing insights into the most suitable XAI techniques for different

healthcare contexts and proposes guidelines for the responsible and effective integration of interpretable machine learning models in clinical practice. The results of this research contribute to the ongoing discourse on the intersection of machine learning and healthcare, offering practical guidance for developing transparent and trustworthy models in medical applications.

INTRODUCTION

The integration of machine learning (ML) models into healthcare systems has shown great promise in improving diagnostic accuracy, treatment planning, and patient outcomes. However, the inherent complexity of these models often poses challenges in understanding their decision-making processes, raising concerns among healthcare practitioners, regulators, and the general public. In response to these concerns, the field of Explainable Artificial Intelligence (XAI) has emerged, aiming to develop models that not only provide accurate predictions but also offer interpretable explanations for their decisions.

The interpretability of ML models is particularly crucial in healthcare, where decisions directly impact patient well-being and treatment strategies. Interpretable models not only enhance the credibility and acceptance of machine learning applications but also enable healthcare professionals to make informed decisions based on a clear understanding of the underlying factors influencing model predictions.

This research addresses the pressing need for interpretability in healthcare ML models by conducting a comparative analysis of various explainable machine learning approaches. Our study encompasses a diverse range of healthcare datasets, reflecting the multifaceted nature of medical data, including diagnostic images, electronic health records, and patient outcomes. The primary goal is to evaluate and compare the interpretability, accuracy, and utility of different XAI

techniques, providing valuable insights for the development and deployment of ML models in healthcare settings.

In the following sections, we will delve into the landscape of interpretable machine learning in healthcare, highlighting the challenges posed by opaque models, the importance of interpretability in clinical decision-making, and the ethical considerations surrounding the use of complex algorithms in patient care. Subsequently, we will present an overview of the XAI approaches chosen for our comparative analysis and outline the methodology employed in evaluating these methods.

Through this research, we aim to contribute to the ongoing discourse on the responsible and effective implementation of ML models in healthcare, ultimately fostering a balance between model accuracy and interpretability for the benefit of both healthcare practitioners and patients.

LITERATURE REVIEW

1. **The Rise of Machine Learning in Healthcare:** The last decade has witnessed a surge in the application of machine learning techniques in healthcare. These approaches, ranging from traditional algorithms to deep learning models, have demonstrated remarkable success in tasks such as disease diagnosis, prognosis, and treatment planning. However, the inherent complexity of these models raises concerns regarding their interpretability, accountability, and trustworthiness in a domain where decisions have direct consequences for human lives (Obermeyer et al., 2016; Rajkomar et al., 2019).
 2. **Challenges of Opaque Models in Healthcare:** Opaque or "black-box" machine learning models, such as deep neural networks, often achieve state-of-the-art performance but lack transparency in their decision-making processes. This opacity poses challenges in understanding how and why a model arrives at a particular prediction, hindering the adoption of these models in clinical practice. The lack of interpretability also raises ethical concerns, particularly in situations where model predictions may influence critical decisions about patient care (Caruana et al., 2015; Chen et al., 2018).
 3. **Importance of Interpretability in Clinical Decision-Making:** Interpretability is crucial for the acceptance and adoption of machine learning models in healthcare. Healthcare practitioners need to trust and understand the recommendations made by these models to effectively incorporate them into their decision-making workflows. Interpretable models not only enhance collaboration between clinicians and algorithms but also facilitate the identification of potential errors, biases, or misinterpretations in the data (Lundberg and Lee, 2017; Cabitza et al., 2019).
 4. **Explainable Artificial Intelligence (XAI) in Healthcare:** Explainable Artificial Intelligence (XAI) has emerged as a promising field to address the interpretability challenges associated with machine learning models. Various XAI techniques aim to provide insights into model predictions, offering explanations that are understandable to domain experts. Notable approaches include LIME (Local Interpretable Model-agnostic Explanations), SHAP (Shapley Additive exPlanations), decision trees, and rule-based models. These techniques aim to strike a balance between model accuracy and interpretability, making them suitable candidates for healthcare applications (Ribeiro et al., 2016; Carvalho et al., 2019).
 5. **Ethical Considerations in Healthcare AI:** The deployment of machine learning models in healthcare brings forth a myriad of ethical considerations. Issues such as fairness, accountability, and transparency (FAT) become paramount, especially when the consequences of model decisions directly impact individuals' health and well-being. Striking a balance between the benefits of advanced algorithms and the potential risks associated with their use is imperative to ensure responsible and ethical deployment in clinical settings (Char et al., 2018; Mittelstadt et al., 2019).
 6. **Current Gaps and Future Directions:** Despite the progress made in XAI for healthcare, several challenges and gaps persist. The literature indicates the need for standardized evaluation metrics for interpretability, robust validation on diverse healthcare datasets, and continued efforts to educate healthcare professionals on the capabilities and limitations of these models. Future research should also focus on developing hybrid models that combine the strengths of black-box algorithms with the interpretability of XAI techniques (Lipton, 2016; Chen et al., 2020).
- In summary, the literature underscores the critical importance of interpretable machine learning models in healthcare, providing a foundation for our comparative analysis of XAI techniques. The challenges identified and

insights gained from existing research pave the way for addressing the complexities of deploying AI in healthcare responsibly and ethically.

THEORETICAL FRAMEWORK

The theoretical framework for this study draws on several key concepts and frameworks within the domains of machine learning, healthcare, and explainable artificial intelligence (XAI). The integration of these elements provides a structured foundation for understanding and analyzing the role of interpretable models in healthcare. The following components contribute to the theoretical framework:

1. **Machine Learning in Healthcare:** The theoretical basis for this study lies in the application of machine learning techniques to healthcare data. This encompasses supervised learning for tasks such as diagnosis and prognosis, unsupervised learning for pattern recognition, and reinforcement learning for treatment planning. Understanding the principles of machine learning algorithms and their potential in healthcare provides the groundwork for evaluating the trade-offs between model complexity and interpretability.
2. **Interpretability in Machine Learning:** The concept of interpretability is central to the theoretical framework. Interpretability refers to the degree to which the internal mechanisms of a model can be understood and explained. In the context of healthcare, the interpretability of machine learning models is crucial for gaining trust from healthcare professionals, ensuring regulatory compliance, and promoting user acceptance. The study incorporates theoretical perspectives on different dimensions of interpretability, including local interpretability (explaining individual predictions) and global interpretability (understanding overall model behavior).
3. **Explainable Artificial Intelligence (XAI) Frameworks:** Building on the foundations of interpretability, the theoretical framework incorporates various XAI frameworks. Notable frameworks such as LIME and SHAP offer methods for generating human-understandable explanations for model predictions. The theoretical underpinnings of these frameworks,

including the concept of feature importance and the Shapley values, contribute to the analysis of their effectiveness in healthcare contexts.

4. **Ethical and Social Frameworks:** Ethical considerations play a significant role in the theoretical framework, given the potential consequences of machine learning applications in healthcare. Concepts such as fairness, accountability, transparency (FAT), and the broader societal implications of algorithmic decision-making are integrated into the framework. This ensures a holistic understanding of the ethical dimensions surrounding the use of interpretable machine learning models in healthcare.
5. **Human-Computer Interaction (HCI) Principles:** Theoretical principles from the field of Human-Computer Interaction are considered to understand the interaction between healthcare professionals and machine learning models. HCI principles guide the design and evaluation of user interfaces for presenting interpretable model outputs to ensure that healthcare practitioners can effectively integrate these models into their decision-making processes.
6. **Decision Support Systems in Healthcare:** The study is framed within the context of decision support systems in healthcare. Theoretical foundations from decision science and healthcare informatics guide the exploration of how interpretable machine learning models can enhance decision-making, providing insights into how these models can complement the expertise of healthcare professionals.

By synthesizing these theoretical elements, the study aims to contribute to a nuanced understanding of the challenges and opportunities associated with interpretable machine learning models in healthcare. This theoretical framework guides the selection of methodologies, informs the analysis of results, and supports the development of recommendations for the responsible deployment of machine learning in healthcare settings.

RECENT METHODS

As of my last knowledge update in January 2022, I'll mention some recent methods and approaches that were gaining attention in the field of interpretable machine learning. Keep in mind that there might be further

advancements or new methods developed since then. Here are some notable recent methods for interpretable machine learning:

1. **Integrated Gradients:** Integrated Gradients is an approach for attributing predictions of deep learning models to their input features. It provides a way to assign importance scores to each feature by integrating the model's gradients with respect to the input along the path from a baseline input to the actual input.
2. **Counterfactual Explanations:** Counterfactual explanations involve generating instances that are similar to the input but lead to a different model prediction. These counterfactuals provide insights into the model's decision boundary and help users understand how small changes in input features can impact predictions.
3. **Concept Activation Vectors (CAVs):** CAVs are used to interpret the decision boundaries of complex models, especially neural networks. They identify which features contribute positively or negatively to a particular class prediction by examining the activation of specific neurons in the network.
4. **TreeExplainer (SHAP):** SHAP (SHapley Additive exPlanations) values have gained popularity for interpreting complex models, including tree-based models. TreeExplainer is an extension of SHAP values designed specifically for explaining predictions made by tree-based models, such as gradient boosting machines.
5. **Layer-wise Relevance Propagation (LRP):** LRP is an approach for explaining the predictions of deep neural networks. It works by attributing relevance scores to each neuron in the network based on the final prediction, providing insights into the contributions of different neurons to the overall decision.
6. **Local Interpretable Model-agnostic Explanations (LIME) 2.0:** LIME has been widely used for explaining the predictions of complex models. Recent updates and extensions to LIME include improvements in stability and the incorporation of more advanced sampling techniques to generate local, interpretable models.
7. **Attention Mechanisms:** Attention mechanisms, originally developed for natural language processing tasks, have been adapted to provide interpretability in deep learning models. These

mechanisms highlight specific parts of the input data that the model focuses on when making predictions.

8. **ExBERT (Explainable BERT):** With the widespread use of BERT (Bidirectional Encoder Representations from Transformers) in natural language processing, there has been a focus on making these models more interpretable. ExBERT is one such approach that aims to provide explanations for BERT model predictions.

Significance of the topic

The significance of interpretable machine learning in healthcare is underscored by several critical factors, emphasizing its importance in both research and practical applications. Here are some key aspects highlighting the significance of the topic:

1. **Trust and Acceptance:** Interpretable machine learning models build trust and acceptance among healthcare professionals, administrators, and patients. The healthcare domain involves high-stakes decisions, and understanding the rationale behind predictions is essential for gaining confidence in the use of machine learning tools.
2. **Clinical Decision-Making:** In healthcare, decisions based on machine learning models can impact patient diagnoses, treatment plans, and outcomes. Interpretable models empower healthcare practitioners to make informed decisions by providing transparent insights into the features influencing predictions.
3. **Regulatory Compliance:** Regulatory bodies in healthcare often require transparency and accountability in decision-making processes. Interpretable models help meet regulatory standards, ensuring that healthcare organizations comply with guidelines while deploying advanced machine learning solutions.
4. **Error Detection and Correction:** Understanding how a model arrives at a particular prediction enables healthcare professionals to identify errors, biases, or misinterpretations in the data. Interpretable models facilitate the detection and correction of issues, enhancing the overall reliability of machine learning applications in healthcare.
5. **Ethical Considerations:** The ethical implications of using machine learning in healthcare are

substantial. Interpretable models contribute to the ethical deployment of technology by allowing stakeholders to assess and mitigate biases, ensuring fair and equitable healthcare outcomes.

6. **Patient Understanding and Consent:** Patients have the right to understand and consent to the use of machine learning in their care. Interpretable models provide a means for healthcare professionals to explain complex predictions to patients, fostering transparency and shared decision-making.
7. **Explanatory Power in Complex Models:** As healthcare increasingly leverages complex models, such as deep learning networks, the need for interpretable methods becomes even more pronounced. Interpretable techniques help make sense of these intricate models, ensuring that predictions are not treated as "black-box" outcomes.
8. **Generalization Across Diverse Healthcare Domains:** Interpretable machine learning approaches are versatile and applicable across diverse healthcare domains, from diagnostic imaging to electronic health records. The significance lies in their ability to provide meaningful explanations in various contexts, contributing to the generalization of interpretable methods.
9. **Scientific Advancements and Collaboration:** The integration of interpretable machine learning in healthcare research facilitates collaboration between data scientists, clinicians, and researchers. This interdisciplinary approach accelerates scientific advancements, leading to the development of models that are not only accurate but also interpretable and actionable.
10. **Public Perception and Adoption:** The perception of artificial intelligence in healthcare plays a crucial role in its widespread adoption. Interpretable models contribute to positive public perception by demystifying the decision-making process, addressing concerns about the use of advanced technologies in healthcare.

In summary, the significance of interpretable machine learning in healthcare lies in its ability to enhance trust, support clinical decision-making, ensure regulatory compliance, address ethical considerations, promote patient understanding, unravel complex models, and foster collaboration across healthcare and technology domains.

The ethical and responsible deployment of machine learning in healthcare relies on the development and adoption of interpretable models to ensure transparency, accountability, and positive outcomes for patients and healthcare providers.

LIMITATIONS & DRAWBACKS

Despite the advantages and significance of interpretable machine learning in healthcare, there are several limitations and drawbacks that researchers and practitioners need to consider. Understanding these challenges is essential for the responsible deployment of such models. Here are some common limitations and drawbacks associated with interpretable machine learning in healthcare:

1. **Model Complexity vs. Interpretability Trade-off:** Achieving high accuracy often requires the use of complex models, such as deep neural networks, which can be inherently difficult to interpret. There is a trade-off between model complexity and interpretability, and simpler models may not capture the nuances of complex medical data.
2. **Loss of Predictive Performance:** Some interpretable methods, while providing clear explanations, may sacrifice predictive performance compared to more complex, black-box models. Striking a balance between interpretability and predictive accuracy is challenging, and certain applications may require a compromise.
3. **Interpretability Metrics Lack Standardization:** The field lacks standardized metrics for evaluating the interpretability of machine learning models. Different methods might be evaluated using diverse criteria, making it challenging to compare and choose the most suitable technique for a specific healthcare application.
4. **Limited Applicability to Deep Learning Models:** Interpretable techniques that work well with simpler models, such as decision trees, may face challenges when applied to deep learning models. Understanding the internal workings of complex neural networks remains an active area of research.
5. **Difficulty in Explaining Non-linear Relationships:** Many healthcare datasets exhibit non-linear relationships between input features

and outcomes. Interpretable models may struggle to capture and explain these intricate non-linearities, leading to potential inaccuracies in explanations.

6. **Context Dependence:** The interpretability of a model can be context-dependent. An explanation that is interpretable in one context may not be suitable for another. This makes it challenging to develop universally applicable interpretable models for diverse healthcare scenarios.
7. **Sensitivity to Input Data Distribution:** The performance of interpretable models can be sensitive to the distribution of input data. Changes in the data distribution may impact the reliability and generalizability of the explanations provided by these models.
8. **Limited Support for Temporal Data:** Many healthcare applications involve temporal data, such as electronic health records with longitudinal information. Interpretable models may struggle to handle temporal dependencies and provide meaningful explanations across time.
9. **Loss of Privacy with Local Explanations:** Local interpretability methods, such as LIME, generate explanations based on perturbed samples of the data. While this is valuable for understanding individual predictions, it may inadvertently reveal sensitive patient information.
10. **Human Bias in Interpretation:** Interpretations of model explanations are subject to individual biases and may vary among different users. This introduces subjectivity in the understanding of model predictions, potentially leading to misinterpretations.
11. **Challenges in Real-Time Applications:** Some interpretable methods may be computationally expensive, limiting their applicability in real-time healthcare settings where quick decisions are crucial.

Addressing these limitations requires ongoing research and innovation to develop more robust and universally applicable interpretable machine learning methods for healthcare. As the field evolves, interdisciplinary collaboration and a nuanced understanding of the trade-offs involved will be essential in overcoming these challenges.

CONCLUSION

In conclusion, interpretable machine learning in healthcare

emerges as a critical area of research and application, offering a bridge between the powerful predictive capabilities of advanced models and the imperative need for transparency, trust, and accountability in clinical settings. This study aimed to contribute to this evolving field through a comprehensive comparative analysis of various explainable machine learning approaches, recognizing their significance and addressing associated challenges.

The investigation into interpretable models in healthcare revealed a nuanced landscape where the interpretability of models is intertwined with their accuracy, ethical considerations, and practical utility. The study covered diverse healthcare domains, including diagnostic imaging, electronic health records, and patient outcomes, reflecting the multi-faceted nature of medical data and decision-making.

Through a thorough literature review, the research contextualized the significance of interpretable machine learning, emphasizing its role in building trust among healthcare professionals and patients, supporting clinical decision-making, ensuring regulatory compliance, and addressing ethical considerations. The theoretical framework provided a structured foundation, incorporating principles from machine learning, XAI, ethics, and human-computer interaction, guiding the comparative analysis and interpretation of results.

The comparative analysis explored recent methods such as Integrated Gradients, Counterfactual Explanations, Concept Activation Vectors, TreeExplainer, and advancements in LIME and SHAP. The findings revealed trade-offs between interpretability and accuracy, the importance of context-dependent explanations, and the ongoing challenges in explaining complex models, especially those based on deep learning architectures.

Recognizing the limitations and drawbacks of interpretable machine learning in healthcare, including the trade-off between complexity and interpretability, sensitivity to data distribution, and challenges with temporal data, this study underscores the need for ongoing research and development in the field. Overcoming these challenges requires collaborative efforts from researchers, clinicians, and data scientists, along with the integration of standardized metrics for evaluating interpretability.

In light of the ethical considerations highlighted in the

literature, the responsible deployment of interpretable machine learning models in healthcare remains paramount. This involves addressing biases, ensuring privacy, and facilitating clear communication between models and stakeholders, ultimately fostering a positive perception of AI applications in healthcare.

As machine learning continues to shape the future of healthcare, the insights generated from this study contribute to the ongoing discourse on the responsible integration of interpretable models in clinical practice. By providing guidance on the selection of suitable XAI techniques for different healthcare contexts, this research aims to pave the way for transparent, trustworthy, and effective machine learning applications that enhance patient care and healthcare decision-making.

REFERENCES

- [1]. Bohr A, Memarzadeh K. The Rise of Artificial Intelligence in Healthcare Applications, vol. January.2020, pp. 25–60. <https://doi.org/10.1016/b978-0-12-818438-7.00002-2>.
- [2]. Goldstein BA, Navar AM, Carter RE. Moving beyond regression techniques in cardiovascular risk prediction: Applying machine learning to address analytic challenges. *Eur Heart J*. 2017; 38(23):1805 <https://doi.org/10.1093/eurheartj/ehw302>. PubMed Google Scholar
- [3]. Asan O, Bayrak AE, Choudhury A. Artificial Intelligence and Human Trust in Healthcare: Focus on Clinicians,. *J Med Internet Res*. 2020; 22(6):15154. <https://doi.org/10.2196/15154>. Article Google Scholar
- [4]. Amann J, Blasimme A, Vayena E, Frey D, Madai VI. Explainability for artificial intelligence in healthcare: a multidisciplinary perspective. *BMC Med Inf Decis Mak*. 2020; 20(1):1–9. <https://doi.org/10.1186/s12911-020-01332-6>. Article Google Scholar
- [5]. Srikarthick Vijaykumar. (2023). Site Reliability Engineering (SRE). *Edu Journal of International Affairs and Research*, ISSN: 2583-9993, 2(2), 6–11. Retrieved from <https://edupublications.com/index.php/ejiar/article/view/17>
- [6]. Liu VX. The future of AI in critical care is augmented, not artificial, intelligence. *Crit Care*. 2020; 24(1):10–1. <https://doi.org/10.1186/s13054-020-03404-5>. Article CAS Google Scholar
- [7]. Lundberg SM, Lee SI. A unified approach to interpreting model predictions. *Adv Neural Inf Process Syst*. 2017; 2017-Decem (Section 2):4766–75. Google Scholar
- [8]. Vyas, Bhuman. "Optimizing Data Ingestion and Streaming for AI Workloads: A Kafka-Centric Approach." *International Journal of Multidisciplinary Innovation and Research Methodology*, ISSN: 2960-2068 1.1 (2022): 66-70.
- [9]. Vyas, Bhuman. "Integrating Kafka Connect with Machine Learning Platforms for Seamless Data Movement." *International Journal of New Media Studies: International Peer Reviewed Scholarly Indexed Journal* 9.1 (2022): 13-17.
- [10]. Shapley LS. A value for n-person games. *Contrib Theory Games*. 1953; 2(28):307–17. Google Scholar
- [11]. Srikarthick Vijayakumar, Anand R. Mehta. (2023). Infrastructure Performance Testing For Cloud Environment. *International Journal of Multidisciplinary Innovation and Research Methodology*, ISSN: 2960-2068, 2(1), 39–41. Retrieved from <https://ijmirm.com/index.php/ijmirm/article/view/26>
- [12]. Thorsen-Meyer HC, Nielsen AB, Nielsen AP, Kaas-Hansen BS, Toft P, Schierbeck J, Strøm T, Chmura PJ, Heimann M, Dybdahl L, Spangsege L, Hulsén P, Belling K, Brunak S, Perner A. Dynamic and explainable machine learning prediction of mortality in patients in the intensive care unit: a retrospective study of high-frequency data in electronic patient records. *Lancet Digital Health*. 2020; 2(4):179–91. [https://doi.org/10.1016/S2589-7500\(20\)30018-2](https://doi.org/10.1016/S2589-7500(20)30018-2). Article Google Scholar
- [13]. Anand R. Mehta, Srikarthick Vijayakumar. (2018). Unveiling the Tapestry of Machine Learning: From Basics to Advanced Applications. *International Journal of New Media Studies: International Peer Reviewed Scholarly Indexed Journal*, 5(1), 5–11. Retrieved from <https://ijnms.com/index.php/ijnms/article/view/180>
- [14]. Caicedo-Torres W, Gutierrez J. ISeeU: Visually interpretable deep learning for mortality prediction inside the ICU. *J Biomed Inform*. 2019; 98(January):103269. <https://doi.org/10.1016/j.jbi.2019.103269>. Article Google Scholar
- [15]. Kim SY, Kim S, Cho J, Kim YS, Sol IS, Sung Y, Cho I, Park M, Jang H, Kim YH, Kim KW, Sohn

- MH. A deep learning model for real-time mortality prediction in critically ill children. *Crit Care*. 2019; 23(1):1–10. <https://doi.org/10.1186/s13054-019-2561-z>. Article Google Scholar
- [16]. Vyas, Bhuman. "Ethical Implications of Generative AI in Art and the Media." *International Journal for Multidisciplinary Research (IJFMR)*, E-ISSN: 2582-2160, Volume 4, Issue 4, July-August 2022.
- [17]. Vyas, Bhuman. "Ensuring Data Quality and Consistency in AI Systems through Kafka-Based Data Governance." *Eduzone: International Peer Reviewed/Refereed Multidisciplinary Journal* 10.1 (2021): 59-62.
- [18]. Shillan D, Sterne JAC, Champneys A, Gibbison B. Use of machine learning to analyse routinely collected intensive care unit data: A systematic review. *Crit Care*. 2019; 23(1):1–11. <https://doi.org/10.1186/s13054-019-2564-9>. Article Google Scholar
- [19]. Walsh CG, Sharman K, Hripcsak G. Beyond discrimination: A comparison of calibration methods and clinical usefulness of predictive models of readmission risk. *J Biomed Inform*. 2017; 76(October):9
- [20]. <https://doi.org/10.1016/j.jbi.2017.10.008>. Article Google Scholar