

Cybersecurity and Privacy Challenges in AI-Driven Healthcare

* **Shamsu Sabo, Isah Muhammad Alhassan, Umar Sani Muhammad**

Faculty of Computing and Sciences, Department of Computer Science, Azman University, Kano, Nigeria.

Correspondence: shamsu.s@azmanuniversity.edu.ng

Abstract

The use of Artificial Intelligence (AI) in healthcare has significantly improved medical diagnostics, patient care, and personalized treatment. However, this progress has introduced important challenges, as AI systems rely on large volumes of sensitive patient data, increasing the risk of cybersecurity breaches and privacy violations. This study adopts a narrative review approach to examine key cybersecurity and data privacy challenges in AI-driven healthcare systems. It explores major threats such as data breaches, adversarial attacks on AI models, and the re-identification of patients from supposedly anonymized data. Ethical concerns related to the collection and use of patient information are also discussed, alongside the strengths and limitations of current mitigation strategies, including federated learning, blockchain technology, differential privacy, and advanced encryption techniques. Drawing on existing literature, the study identifies critical gaps in current security and privacy frameworks and emphasizes the need for stronger regulatory structures, ethical governance, and practical security measures to support the safe and responsible deployment of AI in healthcare.

Keywords: Artificial Intelligence, Cybersecurity, Healthcare Systems, Blockchain, Federated Learning

INTRODUCTION

The rise of artificial intelligence (AI) in healthcare is driving transformative changes in diagnosis, treatment, and patient management. AI systems, powered by machine learning (ML), deep learning, and natural language processing (NLP), excel at analyzing vast health datasets to spot patterns—such as anomalous patient responses or treatment outcomes—that human clinicians might miss. These tools are now widespread in areas such as medical imaging, drug discovery, telemedicine, and clinical decision-support systems (Haug & Drazen, 2023). Yet as AI integrates deeper into healthcare, it expands the attack surface for cybersecurity risks and patient data privacy breaches, threatening data confidentiality,

integrity, and availability (Murdoch, 2021). The rapid expansion of AI adoption in healthcare further increases the urgency of addressing these risks.

Historically, healthcare relied on local, paper-based records with minimal cyber exposure. Rapid digitization, cloud computing, and Internet of Things (IoT) devices have exploded data volumes and vulnerabilities (Affia et al., 2023).

In patient-safety-focused settings, cyber weaknesses can harm clinical outcomes and erode trust. For instance, adversaries might manipulate AI inputs—like altering CT scans to hide lung cancer evidence—leading to misdiagnosis (Mirsky et al., 2019). Similarly, stolen records could enable

identity theft and insurance fraud, damaging patients and institutions. In 2023 alone, U.S. healthcare saw over 700 breaches affecting 133 million records (HHS, 2023). These incidents carry ethical, financial, and social fallout.

AI's distributed training on centralized, often personally identifiable datasets heightens re-identification risks, even in de-identified data (Packhäuser et al., 2022). Regulations like GDPR and HIPAA provide frameworks but lag behind AI's pace, especially with cross-border data flows complicating consent and accountability (Ferrara, 2024). The decentralized and cross-border nature of AI data flows complicates traditional notions of consent, ownership, and accountability.

Mitigation demands a multidisciplinary push: robust cybersecurity, ethical AI governance, and tech like blockchain for tamper-proof audits, federated learning to train models without sharing raw data (Rieke et al., 2020) and privacy enhancers such as homomorphic encryption, differential privacy, and secure multiparty computation (Alnasser & Li, 2025). Drawbacks include high computational demands and expertise gaps in many organizations.

Ethically, insecure or biased datasets can perpetuate discrimination, worsening disparities for underrepresented groups (Hasanzadeh et al., 2025). Transparency, explainability, and accountability are essential for trust and compliance. The World Health Organization (World Health Organization, 2021), frames cybersecurity as a core patient safety issue.

This paper examines cybersecurity and data privacy challenges in AI-driven healthcare systems, evaluates existing mitigation strategies, and identifies a key research gap: the lack of

integrated and empirically validated frameworks for assessing the real-world effectiveness of these approaches, particularly in resource-constrained healthcare environments. It offers policy recommendations for safe, ethical AI adoption.

METHODOLOGY / THREAT ANALYSIS

Literature Search Strategy

A structured literature search was conducted to support the narrative review and threat analysis presented in this study. Between January and April 2025, five major databases—IEEE Xplore, PubMed, ScienceDirect, SpringerLink, and Google Scholar—were used to identify peer-reviewed studies and institutional reports published between 2017 and 2025. Search terms included “AI cybersecurity in healthcare,” “federated learning,” “blockchain health records,” “adversarial attacks on medical AI,” and “privacy-preserving machine learning.”

Inclusion, Exclusion, and Screening Criteria

Inclusion Criteria:

- Peer-reviewed studies examining the application of AI in healthcare and addressing cybersecurity or data privacy concerns.
- References that discuss mitigating threats to privacy or cybersecurity for at least one AI process/system that supports the use of AI in healthcare.
- Peer-reviewed references containing qualitative evaluations, quantitative evaluations, and/or conceptual evaluations of specific threats or mitigations.

Exclusion Criteria:

- Articles that have not been peer-reviewed;

editorial pieces; opinion pieces; or pre-print publications.

- All references must be directly related to healthcare and the use of AI.
- Papers published prior to 2017 and/or non-English language publications; neither are eligible for inclusion in the review process.
- Papers that do not specifically address security and/or privacy issues associated with using AI and/or do not provide mitigation approaches to those issues.

Screening Process:

Through the initial search, 78 potential articles were identified. The researcher subsequently removed any duplicate records, and then all titles and abstracts were reviewed and documented to allow for easy reference while assessing each reference for eligibility. Eligible studies were then reviewed in full text against the inclusion and exclusion criteria. Of the original pool of Articles, only 21 met the inclusion criteria outlined above, and thus only 15 studies with the strongest relevance to healthcare AI security, mitigation strategies, and practical implementation were selected as core references for detailed synthesis. The screening process was conducted carefully to reduce selection bias and ensure consistency in study selection.

Review of Mitigation Technologies

To assess the mitigation technologies (federated learning, blockchain, differential privacy, encryption, and various intrusion detection systems) for this review, the following criteria were used to evaluate each technology:

- Effectiveness of each mitigation technology to mitigate specific threats.
- Impact on the performance/accuracy of each mitigation technology.
- The scalability of each mitigation technology and any computational requirements associated with that technology.
- Feasibility to implement the mitigation technology (and particularly low-resource (LMIC) conditions).
- Level of validation supported by real-world implementation or simulation studies.

RESULTS

Structured Risk Taxonomy

To provide a clearer understanding of the risks, threats were organized into a structured taxonomy based on the extended Confidentiality, Integrity, and Availability (CIA) triad, with additional governance and ethical dimensions. This taxonomy synthesizes findings from the reviewed literature while highlighting new insights relevant to AI-driven healthcare systems.

The different types of threats were evaluated based on their attack surfaces, anticipated probability of occurrence, and severity of consequences it may have on patient safety and the operational readiness of the organization.

Linking Threats to Mitigation Strategies

Another significant shortcoming in the literature is the tendency to address threat types and mitigating strategies separately. In this study, we close this gap by creating a direct link between significant threats and mitigating strategies and evaluating how effective they are in practice.

Table 1: Structured Risk Taxonomy for AI-Driven Healthcare Systems

Threat Category	Description	Primary Impact	Representative Studies
Data Breach & Leakage	Unauthorized access, leakage, or reconstruction of sensitive patient data	Confidentiality & Privacy	Shokri et al. (2017), Murdoch (2021)
Adversarial Attacks	Deliberate manipulation of inputs to mislead AI models (e.g., altered medical images)	Integrity & Clinical Safety	Mirsky et al. (2019), Dong et al. (2024)
Model & Data Integrity Risks	Model inversion, data poisoning, or tampering with training data	Integrity & Trust	Fredrikson et al. (2015), Zhou et al. (2024)
System & Network Attacks	Ransomware, IoMT exploits, and infrastructure-level attacks	Availability & Operations	He et al. (2021), Kingslin & Thasleem (2025)
Governance & Compliance Risks	Failure to comply with GDPR, HIPAA, or ethical standards; consent violations	Regulatory, Ethical & Trust	WHO (2021), Ferrara (2024)

Table 2: Threat-Mitigation Mapping and Analysis

Threat Category	Key Mitigation Strategies	Strengths	Limitations / Gaps	Effectiveness in LMICs
Data Breach & Leakage	Federated Learning + Differential Privacy + Encryption	Prevents raw data sharing, adds noise	Communication overhead, accuracy trade-off	Moderate
Adversarial Attacks	Adversarial training + AI-based IDS	Improves model robustness	High computational cost, limited generalization	Low
Model Inversion & Poisoning	Homomorphic Encryption + Secure Aggregation	Protects model parameters	Significant performance degradation	Low
System & Network Attacks	Blockchain + AI-driven Intrusion Detection	Provides auditability and real-time monitoring	High latency, difficult legacy system integration	Low
Governance & Compliance	Blockchain smart contracts + Policy frameworks	Enhances accountability and traceability	Regulatory lag and enforcement challenges	Very Low

This mapping reveals that while several technical solutions show promise in controlled environments, their integration remains weak, and real-world validation — particularly in resource-constrained settings — is severely limited. These findings informed the development of the proposed Secure AI Healthcare Framework (SAIHF) presented in Section 3.7.

Ethical Risk Dimensions

(Murdoch, 2021) and (World Health Organization, 2021) point out that the violation of personal information privacy in healthcare contexts carries moral considerations that exceed technical harm. Consequently, this study extends the risk model to include ethical risk measures to assess:

- Severity of informed consent violations,
- Risk of bias or discrimination resulting from algorithmic processes,
- Risk of damaging the healthcare institution's reputation.

This dual technical-ethical risk model represents risk in the context of deployed AI medical care holistically.

Comparative Evaluation of Mitigation Strategies

In order to assess the performance of privacy-preserving mechanisms, the study assesses the most prominent defenses aggregated into three categories:

Privacy-Preserving Machine Learning (PPML)

PPML relies on Federated Learning (FL) and Differential Privacy (DP). For example, (Rieke et al., 2020) successfully tested FL of multi-institutional imaging datasets without storing the data in a central location, and (Alnasser & Li, 2025)

added DP to the privacy analysis by reducing the potential that patient data could be re-identified.

At present, the evaluation framework measures the PPML defenses under four components of assessment:

- Accuracy degradation from centralized models,
- Communication cost (MB/epoch),
- Privacy budget to measure the resistance of a leakage, and
- Scaling for systems (hospitals).

All of the studies report a reduction in performance between 5 and 10% when the privacy noise is added, but diagnostic accuracy is still reported as acceptable (Khalid et al., 2023).

Data Governance with Blockchain Technology

The blockchain architecture system acknowledged by (Han et al., 2025; Kasralikar et al., 2025) are assessed on:

- Latency (s) per transaction
- Throughput (TPS)
- Auditability of data
- Regulatory frameworks

The blockchain guarantees data integrity and traceability, however, energy costs and latency (approximately 5–8 seconds for example: on Ethereum-based, the prototype is a significant hindrance to real-time operations in hospitals (Guerra-Manzanares et al., 2023)

AI-Based Intrusion Detection Systems (IDSs)

With the help of hybrid models that combine supervised CNN classifiers with a rule-based anomaly detection model, (Kingslin1 & R2 Iassociate, 2025) explain that intrusion detection within Healthcare 5.0 improvement is facilitated in comparison to using an AI-based model with the rule-based anomaly detection model or an AI model. This approach compares:

- Detection accuracy (% true positives)
- False alarm rate (FAR%)
- Latency in detection (ms)

Current IDS frameworks report detection with accuracy approaching 96%. Nevertheless, these models exhibit issues with generalization, and can only be applied across many hospital network topologies (Murdoch, 2021).

Proposed Integration Framework

Based on insights synthesized from the reviewed literature, this paper recommends a conceptual Secure AI Healthcare Framework (SAIHF) comprising four layers designed to guide future

implementation of secure and privacy-preserving AI in healthcare systems:

- Data Layer – Utilizes federated learning for the training of distributed models and adds differential privacy for model gradient updates.
- Governance Layer – Leverages blockchain contracts for access control and maintains audit records for accountability.
- Detection Layer – Implements an AI-based IDS in order to provide real-time monitoring of abnormal behavior on the network.
- Policy Layer – Coordinates organizational practices with the GDPR, HIPAA, and WHO ethical guidelines.

This multi-layered model reflects best practice guidelines suggested by the (World Health Organization, 2021)and (Rieke et al., 2020), by achieving both accountability and privacy protection.

Evaluation Metrics

The performance of the framework is evaluated according to several multi-dimensional metrics in the following way:

Table 3: Evaluation Metrics and Supporting Studies

Metric Category	Indicators	Measurement
Security Performance	Threat detection rate, false alarm rate	Quantitative (simulation)
Privacy Protection	ϵ -value (DP), model reconstruction success	Quantitative
Operational Efficiency	Latency, throughput, computational overhead	Quantitative
Ethical and Regulatory Compliance	GDPR/HIPAA conformity, explainability, auditability	Qualitative

Simulation studies using available datasets (e.g., MIMIC-III for electronic health records, NIH Chest X-ray for imaging) can validate these indicators using synthetic studies, as examples indicate in (Alnasser & Li, 2025) , as well as (Guerra-Manzanares et al., 2023).

Validation and Limitations

Validation is a conceptual process, not an experimental one, and relies primarily on triangulation from peer-reviewed studies. There are three main limitations to this work:

- **Data Availability** – Real patient datasets are still not available due to privacy legislation.
- **Model Generalizability** – PPML and several of the blockchain-based solutions often encounter limitations in performance when deployed across heterogeneous settings.
- **Interdisciplinary Gaps** – Comprehensive collaboration among clinicians, data scientists, and policy makers is lacking to provide governance.

Future empirical validation would benefit from a partnership with, and deployment of a real-world penetration test within hospitals, of the AI-based clinical systems.

In addition, the integration of multiple technologies within the proposed Secure AI Healthcare Framework (SAIHF) - especially blockchain, federated learning, and AI-based intrusion detection - could also add extra resource utilization and cost. Specifically, incorporating blockchain again means that it will add transaction overhead and high energy consumption. Federated learning requires a lot of communication and processing power, so the singular framework would likely require more resources and be less implementable immediately -

particularly in low resource healthcare settings. Therefore, the proposed framework should be regarded as a conceptual model for future research and evaluation. Further studies should assess its efficacy, cost-effectiveness and scalability prior to implementation.

Ethical Considerations

In accordance with (World Health Organization, 2021) ethical AI guidelines, all of the methods proposed in this work, focus on implementing data minimization, as well as informed consent, fairness, and explainability. According to (World Health Organization, 2021), algorithmic transparency is essential for legal defensibility, and ultimately, public trust. Therefore, the framework should also include components that are 'explainable' AI (XAI) modules to help explain the clinical predictions to the physicians.

DISCUSSION

The narrative provided synthesizes the available current evidence regarding the major challenges of implementing AI-based healthcare systems in relation to cybersecurity and data privacy. An overview of the specific threats, as well as the various mitigation strategies discussed in earlier sections, was presented in order to understand the greater significance; the pragmatic feasibility; and the contextual limitations associated with implementing these strategies.

Implementation Feasibility of Mitigation Strategies

Some privacy-preserving and security technology has been proposed or theorized as having great potential. Theoretical evaluations of federated

learning combined with differential privacy suggest these mechanisms can effectively reduce the risk associated with data leakage. Blockchain technologies provide a method for creating a tamper-proof audit trail for data. However, there are serious barriers to the implementation of these technologies in the real world. Federated learning requires both a sufficient speed (and reliability) of internet connectivity and high computing capacity as the models will continually require updating but neither of these requirements can be satisfied by many current hospitals using these technologies. While providing an effective method for ensuring data provenance through a tamper-proof audit trail, blockchain technologies introduce a long latency (typically 5-8 seconds per transaction) and an extreme amount of energy usage so will not provide a practical solution within the time-sensitive environment of a clinical setting.

Technical Limitations in Low- and Middle-Income Countries (LMICs)

Healthcare organizations in LMICs face significant challenges. Healthcare facilities operating in resource-constrained settings tend to use outdated IT infrastructure, possess limited internet bandwidth, rely on an unreliable power supply, and have a shortage of trained personnel working in cybersecurity. Advanced security techniques (e.g., homomorphic encryption, secure multiparty computation) require immense computational power, typically, the type of power not capable of being accommodated by local servers. AI-Based Intrusion Detection Systems (IDSs) generally need continuous model retraining as well as large annotated datasets; the resources available within LMICs make continuous retraining and reliable sources of annotated datasets scarce.

While proven effective in more developed nations, existing layered security solutions can be expected to decrease significantly in performance and/or be potentially unusable in LMIC contexts without considerable modifications or adaptations. For example, the communication overhead associated with federated learning (FL) becomes excessive when the underlying network bandwidth is in low-bandwidth areas typical of rural/public hospitals.

Key Gaps and Practical Implications

While there is an increased level of interest in layered security solutions; in practice, this interest continues to be expressed through isolated and fragmented solutions. The degree to which researchers have tested empirically integrated frameworks combining FL, blockchain, and IDS has also yet to be experimentally tested within the constraints of actual clinical workflows. The effectiveness of layered security solutions will also be dependent on the level of staff training available to promote successful use as well as on the availability of resources to facilitate change management associated with new technology or process implementation.

Policy and Strategic Recommendations

To develop, the following items should be priorities for healthcare organizations and policymakers:

- More contextually relevant, scalable solutions for low- and middle-income countries, rather than simply adopting models used in wealthier countries.
- Capacity building for both IT staff and clinicians who work with healthcare IT.
- Regulatory sandboxes for real-world testing of integrated AI security frameworks in real-world clinical settings.

- Public-private partnerships to help fill infrastructure and funding gaps.

Practical Considerations and Implications

Although technical approaches like blockchain, federated learning and differential privacy are helpful solutions for many concerns around privacy in the healthcare environment, successful use of

these types of technologies involves more than just technical components. Therefore, health care organizations need to implement a multi-layered approach that includes not only technical safeguards, such as encryption and access control, but also organizational strategies (for example, staff training on data privacy and governance), as well as regulatory compliance (GDPR, HIPAA, WHO Guidelines).

Table 4: Summary of Key Threats, Mitigation Strategies, and Remaining Gaps

Threat Category	Main Mitigation Strategies	Key Remaining Gaps
Data Breach & Leakage	Federated Learning, Encryption, Blockchain audit trails	Limited real-world validation, high implementation cost
Adversarial Attacks	Adversarial training, AI-based IDS	Poor generalization, high computational demands
Model Inversion & Membership Inference	Differential Privacy, Homomorphic Encryption, Secure Aggregation	Performance-privacy trade-offs, scalability issues
System & Network Attacks	AI-driven IDS, Network segmentation, Blockchain	Legacy system integration, latency problems
Governance & Compliance Risks	Smart contracts, Explainable AI (XAI), Policy frameworks	Weak enforcement, especially in LMICs

CONCLUSION

A narrative review of the main types of challenges experienced in AI-based healthcare from a cybersecurity and data privacy perspective has been provided. Cybersecurity issues related to artificial intelligence (AI) in healthcare include data breaches, adversarial attacks (attacks designed to confuse machine learning systems), inversion modeling, and governance problems. The ability to mitigate these threats remains challenging largely due to the high cost of adoption, the complexity of integrating AI

systems into in-use healthcare technologies, lack of real-world data supporting the implementation of given technologies, and the lack of infrastructure or resources in low- and middle-income countries.

POLICY RECOMMENDATIONS

Implementing specific standards for AI security: General guidelines or principles on how to ensure the security of AI systems by addressing the following risks related to the use of AI in healthcare:

adversarial robustness, privacy through not being able to identify a patient based on their behavior captured by the AI on the healthcare systems, etc.

- Implementing privacy-by-design principles: That mandates applying privacy-enhancing technologies to the system design of any healthcare AI system during the development phase.
- Improving cross-sector collaboration: A need to facilitate increased collaboration between the relevant stakeholders which includes health ministries, data protection authorities and AI technology providers to create a cohesive national approach to AI security.
- Increased investment in applied research: Providing needed funding to conduct empirical research and pilot testing integrated security frameworks across the many types of healthcare systems.

Regardless of the technical solutions put in place to secure AI in healthcare, success in this objective requires a long-term commitment. Greater emphasis on contextual adaptation of the systems, rigorous validation of the systems to support their functionality and a commitment to the sustainability of these systems will be required.

REFERENCES

- Affia, A. O., Finch, H., Jung, W., Samori, I. A., Potter, L., & Palmer, X. L. (2023). IoT health devices: Exploring security risks in the connected landscape. *Internet of Things*, 4(2), 150–182. <https://doi.org/10.3390/iot4020009>
- Alnasser, M., & Li, S. (2025). Privacy-enhancing technologies in collaborative healthcare analysis. *Cryptography*, 9(2). <https://doi.org/10.3390/cryptography9020024>
- World Health Organization. (2021). Ethics and governance of artificial intelligence for health: WHO guidance. World Health Organization.
- Ferrara, E. (2024). Fairness and bias in artificial intelligence: A brief survey of sources, impacts, and mitigation strategies. *Sci*, 6(1). <https://doi.org/10.3390/sci6010003>
- Fredrikson, M., Jha, S., & Ristenpart, T. (2015). Model inversion attacks that exploit confidence information and basic countermeasures. *Proceedings of the 22nd ACM SIGSAC Conference on Computer and Communications Security*, 1322–1333. <https://doi.org/10.1145/2810103.2813677>
- Guerra-Manzanares, A., Lopez, L. J. L., Maniatakos, M., & Shamout, F. E. (2023). Privacy-preserving machine learning for healthcare: Open challenges and future perspectives. In *Privacy-Preserving Artificial Intelligence Techniques in Biomedicine and Healthcare* (pp. 51–72). Springer. https://doi.org/10.1007/978-3-031-39539-0_3
- Han, G., Ma, Y., Zhang, Z., & Wang, Y. (2025). A hybrid blockchain-based solution for secure sharing of electronic medical record data. *PeerJ Computer Science*, 11, e2653. <https://doi.org/10.7717/peerj-cs.2653>
- Hasanzadeh, F., Josephson, C. B., Waters, G., Adedinsewo, D., Azizi, Z., & White, J. A. (2025). Bias recognition and mitigation strategies in artificial intelligence healthcare applications. *npj Digital Medicine*, 8(1). <https://doi.org/10.1038/s41746-025-01503-7>
- Haug, C. J., & Drazen, J. M. (2023). Artificial intelligence and machine learning in clinical medicine, 2023. *New England Journal of Medicine*, 388(13), 1201–1208. <https://doi.org/10.1056/NEJMr2302038>
- He, Y., Aliyu, A., Evans, M., & Luo, C. (2021). Health care cybersecurity challenges and solutions under the climate of COVID-19: Scoping review. *Journal of Medical Internet*

- Research, 23(4), e21747.
<https://doi.org/10.2196/21747>
- Jiang, C., Chen, Y., Chang, J., Feng, M., Wang, R., & Yao, J. (2021). Fusion of medical imaging and electronic health records with attention and multi-head mechanisms. *arXiv*.
<https://doi.org/10.48550/arXiv.2112.11710>
- Kasralikar, P., Polu, O. R., Chamarthi, B., Rupavath, R. V. S. S. B., Patel, S., & Tumati, R. (2025). Blockchain for securing AI-driven healthcare systems: A systematic review and future research perspectives. *Cureus*, 17(2).
<https://doi.org/10.7759/cureus.83136>
- Khalid, N., Qayyum, A., Bilal, M., Al-Fuqaha, A., & Qadir, J. (2023). Privacy-preserving artificial intelligence in healthcare: Techniques and applications. *Computers in Biology and Medicine*, 158, 106848.
<https://doi.org/10.1016/j.compbimed.2023.106848>
- Kingslin, S., & Thasleem, T. (2025). AI-driven threat intelligence in healthcare cybersecurity: A comprehensive survey. *International Journal of Research and Scientific Innovation*.
<https://doi.org/10.51244/IJRSI>
- Mirsky, Y., Mahler, T., Shelef, I., & Elovici, Y. (2019). CT-GAN: Malicious tampering of 3D medical imagery using deep learning. *arXiv*. <http://arxiv.org/abs/1901.03597>
- Murdoch, B. (2021). Privacy and artificial intelligence: Challenges for protecting health information in a new era. *BMC Medical Ethics*, 22(1).
<https://doi.org/10.1186/s12910-021-00687-3>
- Packhäuser, K., Gündel, S., Münster, N., Syben, C., Christlein, V., & Maier, A. (2022). Deep learning-based patient re-identification is able to exploit the biometric nature of medical chest X-ray data. *Scientific Reports*, 12(1).
<https://doi.org/10.1038/s41598-022-19045-3>
- Rieke, N., Hancox, J., Li, W., Milletari, F., Roth, H. R., Albarqouni, S., Bakas, S., Galtier, M. N., Landman, B. A., Maier-Hein, K., Ourselin, S., Sheller, M., Summers, R. M., Trask, A., Xu, D., Baust, M., & Cardoso, M. J. (2020). The future of digital health with federated learning. *npj Digital Medicine*, 3(1).
<https://doi.org/10.1038/s41746-020-00323-1>
- Shokri, R., Stronati, M., Song, C., & Shmatikov, V. (2017). Membership inference attacks against machine learning models. *Proceedings of the IEEE Symposium on Security and Privacy*, 3–18.
<https://doi.org/10.1109/SP.2017.41>
- Siqueira, L. P., Batista, C. L., Lui, P. H., Kazienko, J. F., Quincozes, S. E., Quincozes, V. E., Welfer, D., & Nomura, S. (2025). A comprehensive survey on intrusion detection systems for Healthcare 5.0: Concepts, challenges, and practical applications. *Sensors*, 25(20), 6261.
<https://doi.org/10.3390/s25206261>
- Sartor, G., & Lagioia, F. (2020). The impact of the General Data Protection Regulation (GDPR) on artificial intelligence. *European Parliamentary Research Service*.
<https://doi.org/10.2861/293>
- Zhou, Z., Zhu, J., Yu, F., Li, X., Peng, X., Liu, T., & Han, B. (2024). Model inversion attacks: A survey of approaches and countermeasures. *arXiv*.
<http://arxiv.org/abs/2411.10023>