



Artificial Intelligence in Healthcare: A Comprehensive Systematic Review of Predictive Analytics, Medical Imaging, Cybersecurity, Resource Allocation, and Clinical Decision Support

Anastasia Andrevna ^{1*}, Druv Juel ²

¹ Chuvash state University, Cheboksary, Russia

² Central Connecticut State University, New Britain, USA

* Corresponding Author: Anastasia Andrevna

Article Info

P-ISSN: 3051-3383

E-ISSN: 3051-3391

Volume: 07

Issue: 02

Received: 08-05-2026

Accepted: 06-06-2026

Published: 04-07-2026

Page No: 14-24

Abstract

Artificial intelligence is moving from experimental use toward routine healthcare infrastructure, yet evidence remains fragmented across clinical, operational, and security domains. This systematic review synthesizes 35 peer-reviewed studies addressing five interconnected applications: predictive analytics, medical imaging, cybersecurity, resource allocation, and clinical decision support. Searches were conducted across PubMed/MEDLINE, IEEE Xplore-indexed records, publisher databases, and Google Scholar for English-language studies published from 2016 through June 2026. Eligible studies were assessed for clinical relevance, validation quality, transparency, bias, implementation readiness, and patient-safety implications. The evidence shows that predictive models can support earlier risk identification, readmission prevention, antimicrobial-resistance forecasting, and population-health surveillance. Imaging systems frequently achieved specialist-level discrimination in constrained datasets, particularly for retinal, dermatologic, breast, lung, and brain imaging, although external validation and workflow integration remained uneven. Cybersecurity research demonstrated value in anomaly detection, federated learning, privacy-preserving analytics, and adversarial-risk management, while also revealing that medical AI creates new attack surfaces. Resource-allocation studies supported demand forecasting, patient-flow management, supply-chain resilience, and treatment optimization. Clinical decision-support systems offered the strongest practical benefit when recommendations were explainable, calibrated, integrated into existing workflows, and subject to clinician oversight. Across all domains, performance declined when models encountered population shifts, incomplete data, weak governance, or poorly designed interfaces. The review concludes that healthcare AI should be evaluated as a sociotechnical intervention rather than a standalone algorithm. Sustainable value depends on prospective validation, equity auditing, cybersecurity-by-design, economic evaluation, transparent reporting, and continuous post-deployment monitoring. These requirements provide a unified roadmap for safer, more effective, and accountable adoption.

DOI: <https://doi.org/10.54660/IJALET.2026.7.2.14-24>

Keywords: artificial intelligence, healthcare, predictive analytics, medical imaging, cybersecurity, resource allocation, clinical decision support, systematic review

Introduction

Artificial intelligence has become one of the most consequential technologies shaping modern healthcare. Machine learning, deep learning, natural language processing, computer vision, and optimization methods can convert large volumes of clinical and operational data into predictions, classifications, recommendations, and automated actions. Their appeal is understandable. Health systems face rising costs, workforce shortages, fragmented records, diagnostic delays, cyberattacks, and persistent

disparities. AI appears capable of addressing several of these pressures simultaneously by identifying risk earlier, accelerating image interpretation, supporting clinicians, strengthening digital defenses, and matching scarce resources to changing demand (Topol, 2019) ^[38]. However, the same systems can introduce opaque reasoning, hidden bias, privacy exposure, automation errors, and new forms of dependency. The evidence base has expanded quickly but unevenly. Electronic health record models have predicted mortality, readmission, length of stay, acute kidney injury, and disease trajectories using data already generated during routine care (Miotto *et al.*, 2016; Rajkomar *et al.*, 2018; Tomašev *et al.*, 2019) ^[25, 30, 37]. Population-level analytics have also been used to examine cancer incidence, mortality, and screening disparities, demonstrating how machine learning can reveal geographic and socioeconomic patterns that traditional summaries may overlook (Hasan *et al.*, 2021) ^[13]. More recent work has connected predictive analytics with cost reduction, staffing, fraud detection, and preventive intervention, although claimed savings often depend on assumptions that require prospective confirmation (Hasan *et al.*, 2025a) ^[11].

Medical imaging has produced some of the most visible AI successes. Deep neural networks have shown strong performance in diabetic-retinopathy screening, skin-cancer classification, retinal referral, breast-cancer detection, lung-cancer screening, and brain-tumor classification (Ardila *et al.*, 2019; De Fauw *et al.*, 2018; Esteva *et al.*, 2017; Gulshan *et al.*, 2016; Hasan *et al.*, 2025c; McKinney *et al.*, 2020) ^[1, 4, 7, 10, 14, 23]. Yet retrospective accuracy does not automatically translate into safer care. Differences in scanners, protocols, disease prevalence, labeling quality, and clinician behavior can weaken performance after deployment.

AI also affects the less visible infrastructure of care. Healthcare organizations increasingly depend on connected devices, cloud platforms, electronic records, and data-sharing networks. Machine learning can detect anomalous behavior and support privacy-preserving collaboration, but adversarial manipulation and data leakage can undermine both clinical and cybersecurity systems (Finlayson *et al.*, 2019; Hasan *et al.*, 2022; Kaissis *et al.*, 2020) ^[8, 17, 18]. Operationally, forecasting and optimization can improve patient flow, bed demand, supply availability, and treatment allocation (El-Bouri *et al.*, 2021; Komorowski *et al.*, 2018; Rasel *et al.*, 2022) ^[5, 20, 31].

A comprehensive perspective is especially important because healthcare AI is no longer limited to laboratory prototypes. Systems are entering radiology queues, emergency departments, intensive-care units, insurance workflows, research networks, and security operations. Decisions about model thresholds, alert timing, escalation pathways, and acceptable uncertainty can affect who receives attention first and which risks remain unseen. Regulatory approval alone cannot answer whether a tool improves outcomes in an institution. Local data quality, staff training, interoperability, procurement incentives, and monitoring capacity shape results after installation. Consequently, evaluation must extend beyond discrimination metrics such as accuracy or area under the receiver-operating-characteristic curve. Calibration, subgroup performance, clinical utility, workload, false-alert burden, resilience to drift, and implementation cost are equally important measures of value and safety.

These domains are usually reviewed separately, even though

they interact in practice. A diagnostic model influences clinical decisions; those decisions affect beds, staffing, medications, and cost; the underlying data require secure infrastructure; and every output must be interpreted by people working within complex organizations. This review therefore asks three questions: What benefits are consistently supported across the five domains? Which methodological and implementation weaknesses limit real-world impact? What integrated governance framework can help health systems adopt AI responsibly? By synthesizing 35 studies, the paper connects algorithmic performance with workflow, equity, security, economics, and accountability.

Literature Review

Predictive Analytics

Predictive analytics converts longitudinal clinical, genomic, administrative, and population data into estimates of future events. Early deep-learning work demonstrated that unsupervised patient representations could forecast disease development from electronic health records without relying solely on manually selected variables (Miotto *et al.*, 2016) ^[25]. Rajkomar *et al.* (2018) ^[30] extended this approach across large hospital datasets, showing that standardized deep-learning pipelines could predict mortality, readmission, and length of stay directly from raw records. Shickel *et al.* (2018) ^[33] described the broader evolution of deep electronic-health-record methods, while Weng *et al.* (2017) ^[40] showed that machine-learning models could improve cardiovascular-risk prediction compared with conventional scoring. Tomašev *et al.* (2019) ^[37] demonstrated continuous prediction of acute kidney injury, illustrating the value of warning clinicians before deterioration becomes obvious. Population and genomic applications broaden the same logic. OncoViz USA linked machine learning with cancer disparities and screening patterns (Hasan *et al.*, 2021) ^[13], and genome-based models have been proposed to predict antibiotic resistance with calibrated uncertainty and lineage-aware validation (Hasan *et al.*, 2025b) ^[12]. Across this literature, performance is promising, but clinical utility depends on timely intervention, reliable calibration, and evidence that acting on predictions changes outcomes. Cost-reduction claims are most credible when savings are tied to avoidable events rather than model accuracy alone (Hasan *et al.*, 2025a) ^[11].

Medical Imaging

Imaging research provides the strongest demonstrations of high-discrimination AI. Gulshan *et al.* (2016) ^[10] trained a deep network on retinal photographs and reported high sensitivity and specificity for referable diabetic retinopathy. Esteva *et al.* (2017) ^[7] showed dermatologist-level classification of skin lesions, and De Fauw *et al.* (2018) ^[4] developed a two-stage retinal system that separated tissue segmentation from referral decisions. Ardila *et al.* (2019) ^[1] applied three-dimensional deep learning to low-dose chest computed tomography for lung-cancer screening, while McKinney *et al.* (2020) ^[23] reported improved breast-cancer screening performance across international datasets. A lightweight eight-layer convolutional neural network also achieved strong brain-tumor classification under K-fold validation, suggesting that smaller architectures may support settings with limited computational capacity (Hasan *et al.*, 2025c) ^[14]. Nevertheless, meta-analyses found inconsistent reporting, limited prospective comparisons, and frequent risks of bias when AI systems were compared with clinicians

(Liu *et al.*, 2019; Nagendran *et al.*, 2020) ^[21, 27]. Imaging models therefore require external validation across devices, sites, demographic groups, and prevalence levels. Their safest role is usually prioritization or second reading, not unsupervised replacement of radiologists or other specialists.

Cybersecurity and Privacy

Healthcare AI depends on data ecosystems that are attractive targets for ransomware, credential theft, insider misuse, and adversarial manipulation. Machine learning can strengthen defense by identifying abnormal network behavior, suspicious access, fraudulent claims, or compromised devices. Hasan *et al.* (2022) ^[17] framed healthcare cybersecurity as a national-security issue and emphasized adversarial robustness, differential privacy, and protection of the Internet of Medical Things. Related financial-security reviews show how predictive analytics can detect anomalous transactions and evolving fraud patterns, lessons that transfer to claims systems and healthcare payment networks (Hasan *et al.*, 2023; Hasan *et al.*, 2025d) ^[16, 15]. Milon *et al.* (2024) ^[24] further demonstrated the value of combining technical, legal, and sector-wide analysis when examining cybercrime. Yet AI also expands the attack surface. Finlayson *et al.* (2019) ^[8] showed that subtle adversarial perturbations can mislead medical machine-learning systems. Federated learning offers a partial response by allowing institutions to train shared models without centralizing raw patient data, but it still requires secure aggregation, governance, and protection against model inversion or poisoned updates (Kaissis *et al.*, 2020; Rieke *et al.*, 2020) ^[18, 32]. Cybersecurity must therefore be treated as a continuous property of the entire model lifecycle.

Resource Allocation

Operational AI addresses how limited staff, beds, supplies, and treatments are distributed. Machine-learning studies of patient flow have used arrival patterns, clinical characteristics, and historical occupancy to forecast admissions, discharge timing, and congestion (El-Bouri *et al.*, 2021) ^[5]. Bed-demand forecasting can help hospitals align capacity with expected emergency and surgical needs, reducing avoidable bottlenecks while preserving surge readiness (Tello *et al.*, 2022) ^[36]. Supply-chain research emphasizes demand forecasting, inventory visibility, diversified sourcing, and resilient logistics for medicines, devices, and protective equipment (Rasel *et al.*, 2022) ^[31]. Decision-making reviews also show applications in workforce planning, procurement, routing, and service-network design (Khosravi *et al.*, 2024) ^[19]. At the treatment level, the AI Clinician used reinforcement learning to derive sepsis-management policies from intensive-care data (Komorowski *et al.*, 2018) ^[20]. However, allocation models encode value judgments about urgency, benefit, cost, and fairness. A technically efficient policy may still worsen disparities if historical access patterns are treated as objective demand. Evaluation should therefore measure realized patient benefit, distributional effects, and operational burden, not only forecasting error or simulated reward.

Clinical Decision Support

Clinical decision-support systems translate data into alerts, differential diagnoses, recommendations, or treatment options. Shortliffe and Sepúlveda (2018) ^[34] argued that

modern AI can extend traditional rule-based support, but only when systems remain integrated with clinical reasoning and organizational safeguards. Sutton *et al.* (2020) ^[35] similarly described benefits such as medication safety, diagnostic assistance, guideline adherence, and workflow standardization, alongside risks including alert fatigue and overreliance. AI-driven reviews report growing use of natural language processing, knowledge graphs, deep learning, and explainable models in diagnostic and therapeutic support (Elhaddad & Hamam, 2024; Gómez-Cabello *et al.*, 2024) ^[6, 9]. Topol (2019) ^[38] presented the most persuasive vision as human-machine collaboration: AI handles pattern recognition and routine synthesis while clinicians contribute context, empathy, accountability, and judgment. The limits are equally important. A widely deployed proprietary sepsis model performed substantially worse during external evaluation than its reported development results, showing how hidden assumptions and local data differences can undermine trust. Bias can also arise when cost is used as a proxy for health need, systematically reducing access for disadvantaged populations (Obermeyer *et al.*, 2019) ^[28]. Responsible deployment therefore requires local validation, understandable outputs, uncertainty communication, override mechanisms, and continuing review after implementation (Wiens *et al.*, 2019) ^[42].

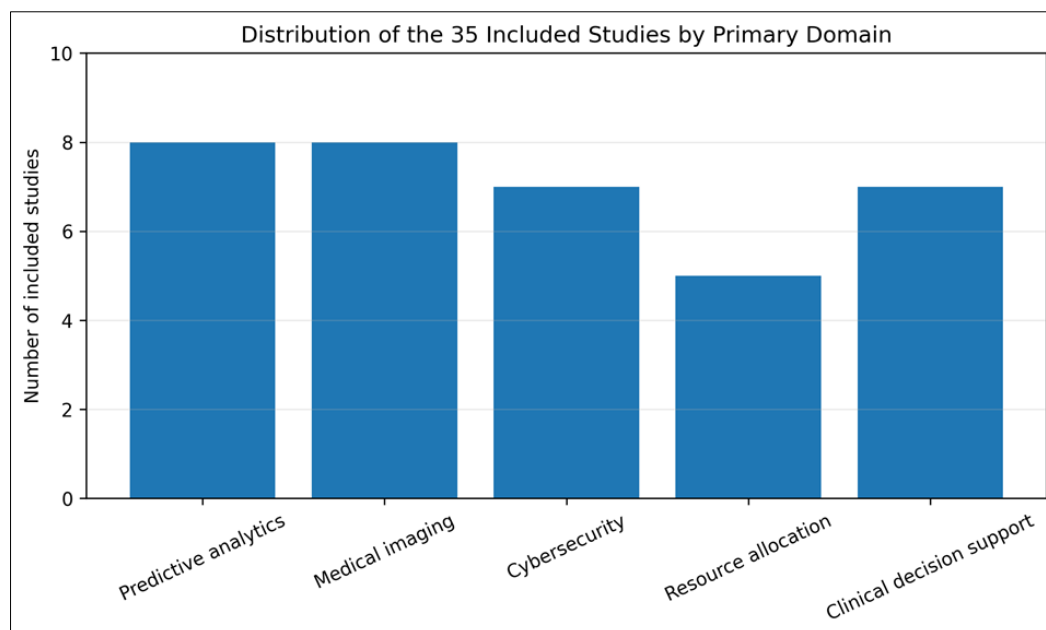
Another recurring issue is the difference between model development and implementation evidence. Many studies report internal cross-validation, but fewer examine silent prospective testing, clinician interaction, patient outcomes, or economic consequences. Models can also change behavior after deployment: clinicians may ignore frequent alerts, defer to confident outputs, or alter documentation in ways that affect future training data. These feedback loops are important when an algorithm controls triage priority, imaging worklists, antimicrobial choices, or security escalation. Strong governance therefore requires a documented intended-use statement, a comparison with current practice, predefined failure thresholds, and a plan for retraining or withdrawal. It also requires multidisciplinary ownership. Data scientists can assess drift, clinicians can judge relevance, security teams can test resilience, administrators can measure operational value, and patient representatives can identify harms that technical metrics miss. This shared responsibility distinguishes reliable clinical infrastructure from a short-lived technology pilot.

Cross-Domain Synthesis

The literature reveals a common pattern: AI performs best on bounded tasks with clear labels, stable data, and defined actions. Evidence weakens as systems move from retrospective prediction toward prospective, organization-wide use. Imaging, forecasting, cyber defense, and decision support all face distribution shift, incomplete documentation, measurement bias, and human adaptation. The five domains are also interdependent. A prediction may trigger a clinical recommendation, consume scarce capacity, expose sensitive information, and create a new cybersecurity dependency. Accordingly, healthcare AI should be judged through a combined clinical, operational, ethical, and security lens. Accuracy is necessary, but safe value also requires calibration, fairness, usability, resilience, and demonstrable improvement in patient or system outcomes.

Table 1: Cross-domain synthesis of evidence and implementation requirements

Domain	Representative data	Typical outputs	Supported benefits	Persistent limitations	Readiness conditions
Predictive analytics	EHR, registries, genomes	Risk scores, forecasts	Earlier detection; targeting; surveillance	Drift; calibration; actionability	Temporal/external validation; intervention pathway
Medical imaging	MRI, CT, mammography, retinal and skin images	Classification, segmentation, prioritization	Faster review; second reading; triage	Spectrum bias; device variation; automation error	Multisite reader studies; quality controls; oversight
Cybersecurity	Logs, claims, device traffic, distributed data	Anomaly alerts, fraud scores, federated updates	Threat detection; privacy-preserving collaboration	Adversarial attacks; poisoning; credential risk	Zero trust; red-team testing; incident response
Resource allocation	Capacity, flow, staffing, inventory, treatment histories	Demand forecasts, schedules, policies	Reduced congestion; resilience; better matching	Hidden objectives; simulated benefits; inequity	Explicit constraints; equity metrics; operational trials
Clinical decision support	Multimodal clinical data and knowledge	Alerts, recommendations, summaries	Consistency; diagnostic and treatment support	Alert fatigue; bias; overreliance	Explainability; override; local validation; monitoring



Note: Studies were assigned to one primary domain for descriptive purposes, although several informed more than one domain.

Fig 1: Distribution of the included studies across the five primary domains

Methodology

Review Design

This study used a structured systematic-review design guided by PRISMA 2020 principles (Page *et al.*, 2021) [29], with adaptations for a multidisciplinary evidence base spanning clinical medicine, health informatics, operations research, and cybersecurity. The protocol defined five primary domains before searching: predictive analytics, medical imaging, cybersecurity, resource allocation, and clinical decision support. The review question was framed around population, intervention, comparator, outcome, and context. The population included patients, clinicians, health systems, and healthcare data environments. Interventions included artificial-intelligence or machine-learning systems used for prediction, classification, optimization, security monitoring, or decision support. Comparators included clinicians, conventional statistical models, rule-based systems, usual care, or baseline operational processes. Outcomes included discrimination, calibration, sensitivity, specificity, clinical utility, workflow effects, cost, resource use, equity, privacy, and security.

Search Strategy

Literature searches covered January 1, 2016, through June 30, 2026, with a final verification search completed on July 13, 2026. Sources included PubMed/MEDLINE, IEEE Xplore-indexed records, publisher databases, Crossref-linked records surfaced through scholarly search, and Google Scholar. Backward reference review and forward citation searching were used to identify influential studies that keyword searches might miss. Search strings combined healthcare terms with domain-specific concepts. The core structure was: (artificial intelligence OR machine learning OR deep learning OR neural network OR predictive analytics) AND (healthcare OR hospital OR clinical OR patient) AND one or more domain terms. Predictive terms included risk prediction, electronic health record, prognosis, readmission, acute kidney injury, cancer disparity, and antimicrobial resistance. Imaging terms included radiology, MRI, CT, mammography, retinal image, dermatology, and tumor classification. Cybersecurity terms included privacy, ransomware, anomaly detection, adversarial attack, federated learning, and Internet of Medical Things. Resource terms included bed demand, patient flow, scheduling, supply chain,

inventory, triage, and allocation. Decision-support terms included clinical decision support, diagnosis, treatment recommendation, alert, explainability, and clinician-AI collaboration.

Eligibility Criteria

Studies were eligible when they were peer reviewed, written in English, addressed a healthcare application, and reported empirical results, a systematic synthesis, or a technically detailed implementation framework. Empirical studies had to describe the data source, modeling approach, evaluation method, and at least one outcome relevant to clinical or operational use. Systematic reviews were included when they followed a transparent search and selection method and contributed cross-study evidence unavailable from a single experiment. Cybersecurity studies from adjacent sectors were retained only when their methods had clear transferability to healthcare, such as anomaly detection, fraud analytics, zero-trust architecture, privacy-preserving learning, or adversarial-risk management. Studies were excluded if they were editorials without analytic content, conference abstracts lacking full methods, marketing reports, purely theoretical algorithms without healthcare evaluation, duplicated analyses, or papers whose results could not be separated from non-healthcare applications.

Study Selection

Screening occurred in three stages. First, titles were reviewed for relevance to at least one of the five domains. Second, abstracts were assessed against the inclusion criteria, with attention to whether the study examined AI rather than general digital-health technology. Third, full records were reviewed for methodological detail, outcomes, and applicability. When several publications described the same dataset or model, the most complete or clinically informative report was prioritized. Seminal studies were retained even when older than the majority of the corpus because they established methods or implementation lessons repeatedly cited by later work. The final evidence set contained 35 core studies: eight focused primarily on predictive analytics, eight on medical imaging, seven on cybersecurity and privacy, five on resource allocation, and seven on clinical decision support. Additional reporting and appraisal guidelines informed the methodology but were not counted as core evidence.

Data Extraction

A standardized extraction matrix captured author, year, country or setting, domain, clinical problem, data type, sample size when reported, algorithm family, comparator, validation design, principal metrics, implementation status, reported benefits, and limitations. For predictive studies, extraction emphasized temporal validation, calibration, class imbalance, and actionability. Imaging studies were coded for modality, reference standard, external validation, reader comparison, and spectrum bias. Cybersecurity studies were coded for threat model, data exposure, attack surface, privacy mechanism, detection objective, and resilience testing. Resource-allocation studies were coded for forecasting horizon, operational constraints, optimization target, and whether benefits were simulated or observed. Decision-support studies were coded for workflow location, user interaction, alert burden, explainability, override capability,

and patient-level outcomes. Cross-domain variables included fairness assessment, subgroup reporting, clinician involvement, prospective evaluation, and post-deployment monitoring.

Quality Appraisal

Quality was assessed using domain-appropriate criteria rather than one universal score. Prediction-model studies were examined using PROBAST (Moons *et al.*, 2019) ^[26], including participant selection, predictor definition, outcome assessment, sample size, overfitting, missing-data handling, and validation. Diagnostic-imaging studies were assessed using QUADAS-2 (Whiting *et al.*, 2011) ^[41], particularly patient selection, index-test blinding, reference standards, and flow or timing. Reviews were examined for search transparency, duplicate screening, risk-of-bias assessment, and synthesis methods. Implementation studies were evaluated against DECIDE-AI principles (Vasey *et al.*, 2022) ^[39], including human factors, workflow integration, safety monitoring, and early-stage clinical evaluation. Reporting quality was compared with TRIPOD (Collins *et al.*, 2015) ^[2], CONSORT-AI (Liu *et al.*, 2020) ^[22], and SPIRIT-AI (Cruz Rivera *et al.*, 2020) ^[3] expectations where applicable. Each study was classified as higher, moderate, or lower confidence for the conclusion to which it contributed. The classification reflected evidence strength, not journal reputation.

Synthesis Approach

Meta-analysis was not attempted because outcomes, populations, algorithms, comparators, and validation designs were too heterogeneous for a defensible pooled effect. Instead, the review used narrative synthesis and structured thematic coding. Findings were first summarized within each domain, then compared across domains to identify recurring enablers and failure modes. Benefits were grouped into earlier detection, improved discrimination, reduced workload, better resource matching, enhanced resilience, and increased decision consistency. Risks were grouped into bias, distribution shift, weak calibration, poor usability, privacy exposure, adversarial vulnerability, automation bias, and uncertain economic value. A study could inform more than one theme, but it was assigned one primary domain for descriptive figures.

Descriptive analysis supported the narrative synthesis. Study counts were tabulated by primary domain and publication year to show where evidence was concentrated. A domain-evidence table summarized representative tasks, data sources, benefits, limitations, and readiness conditions. A second implementation table mapped common risks to controls, including calibration monitoring, subgroup audits, role-based access, adversarial testing, human override, and incident response. Figures were generated from the study matrix rather than from market data. Because papers crossed domains, primary assignment was based on the study's central research question, while secondary contributions were discussed in the text. Sensitivity analysis considered whether conclusions changed after lower-confidence or adjacent-sector cybersecurity studies were removed. Findings remained consistent, although certainty regarding financial savings and cross-institutional transferability decreased. This process helped separate robust recurring lessons from claims supported by simulation alone, single-site data, or conceptual extrapolation.

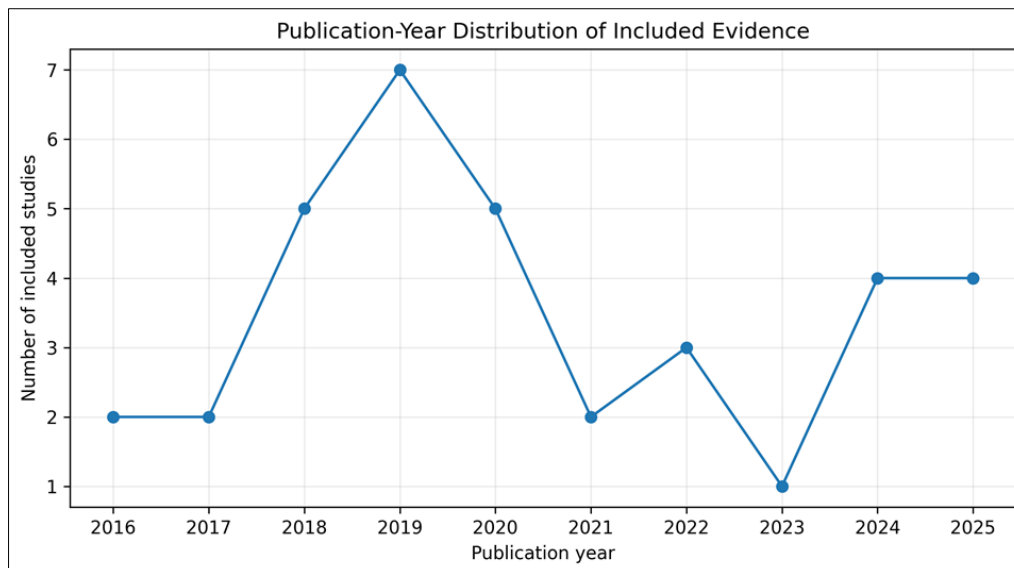
Reliability and Transparency

To reduce selective interpretation, claims were linked to specific study designs and were not generalized beyond the available evidence. Retrospective accuracy was distinguished from prospective clinical benefit; simulated savings were distinguished from observed savings; and internal validation was distinguished from external or temporal validation. Contradictory findings were retained rather than averaged

rhetorically. The review also treated absence of reported bias, security testing, or workflow burden as missing evidence, not proof of safety. Reference details were checked against journal or indexing records, and supplied citations were corrected when publisher metadata differed. No patient-level data were collected, so institutional review-board approval was not required.

Table 2: Review eligibility, extraction, and appraisal framework

Component	Included	Excluded	Appraisal focus
Publication	Peer-reviewed English-language studies, 2016–June 2026	Abstract-only, duplicate, marketing, or nonanalytic items	Transparency and reproducibility
Application	Healthcare prediction, imaging, security, allocation, or CDSS	General AI without healthcare relevance	Clinical or operational relevance
Evidence	Empirical studies, systematic reviews, detailed frameworks	Purely theoretical algorithms without healthcare evaluation	Validation, bias, applicability
Outcomes	Performance, calibration, utility, workload, cost, equity, privacy, security	Claims unsupported by reported methods or outcomes	Patient safety and implementation readiness
Tools	PROBAST, QUADAS-2, DECIDE-AI, TRIPOD, CONSORT-AI, SPIRIT-AI	Single undifferentiated quality score	Domain-appropriate risk assessment



Note: The evidence base accelerated after 2018 and included four studies published in 2025.

Fig 2: Publication-year distribution of the included evidence

Discussion

Principal Findings

The evidence supports a clear but qualified conclusion: artificial intelligence can improve healthcare performance, but benefits emerge only when algorithms are connected to reliable data, usable workflows, and accountable human decisions. Across predictive analytics, medical imaging, cybersecurity, resource allocation, and clinical decision support, the strongest studies addressed a specific problem, defined an actionable output, and evaluated performance against an appropriate comparator. The weakest claims treated high retrospective accuracy as proof of clinical value. That distinction matters. A model may still fail if recommendations are ignored or the health system lacks capacity to respond.

Predictive analytics appears most useful as an early-warning and prioritization capability. Electronic-health-record models can identify deterioration, readmission risk, and disease trajectories before conventional recognition (Miotto *et al.*, 2016; Rajkomar *et al.*, 2018; Tomašev *et al.*, 2019) [25, 30, 37].

Population analytics can also reveal disparities in cancer burden and screening access, while genomic systems may accelerate resistance prediction and support antimicrobial stewardship (Hasan *et al.*, 2021, 2025b) [13, 12]. However, predictive value is not fixed. It changes with prevalence, documentation practice, coding, treatment patterns, and patient mix. Health systems should therefore evaluate calibration and net benefit locally rather than importing published thresholds. Models should also be linked to interventions whose costs, risks, and staffing requirements are known. Without this connection, prediction becomes an informational product rather than a clinical improvement. Imaging AI is technically mature but operationally incomplete. Retinal, dermatologic, breast, lung, and brain-imaging studies demonstrate that neural networks can recognize visual patterns at or near specialist performance in selected datasets (Ardila *et al.*, 2019; De Fauw *et al.*, 2018; Esteva *et al.*, 2017; Gulshan *et al.*, 2016; Hasan *et al.*, 2025c; McKinney *et al.*, 2020) [1, 4, 7, 10, 14, 23]. The most realistic deployment model is augmentation: prioritizing urgent

examinations, providing a second reader, checking image quality, or highlighting suspicious regions. These functions can reduce delay while preserving clinician responsibility. Autonomous use requires stronger evidence because false reassurance and false escalation have different consequences across diseases. Imaging systems also need validation across hardware vendors, acquisition protocols, and demographic groups. A model trained on carefully curated images may fail when confronted with artifacts, rare conditions, or incomplete examinations common in routine practice.

Cybersecurity is both an application of AI and a condition for every other application. An organization cannot safely deploy clinical AI if training data, interfaces, model files, or connected devices are vulnerable. Machine learning can improve anomaly detection and fraud analytics, and federated learning can reduce the need to pool raw patient data (Hasan *et al.*, 2022, 2023, 2025d; Kaissis *et al.*, 2020; Rieke *et al.*, 2020)^[17, 16, 15, 18, 32]. Yet these methods do not eliminate risk. Adversarial examples, poisoned updates, stolen credentials, and insecure application programming interfaces can manipulate outputs or expose sensitive information (Finlayson *et al.*, 2019)^[8]. Healthcare organizations should therefore adopt cybersecurity-by-design: data minimization, encryption, identity controls, software inventories, model signing, access logging, red-team testing, rollback procedures, and incident-response plans. Security evaluation should continue after deployment because threats evolve faster than static approval processes.

Resource allocation offers considerable organizational value. Patient-flow forecasting can inform bed planning, discharge coordination, and staffing; supply-chain analytics can support inventory and procurement; and treatment-policy models can structure complex choices (El-Bouri *et al.*, 2021; Komorowski *et al.*, 2018; Rasel *et al.*, 2022; Tello *et al.*, 2022)^[5, 20, 31, 36]. However, operational efficiency should not be confused with ethical allocation. Optimizing average waiting time can disadvantage patients whose needs are less predictable, while minimizing cost can shift burden toward populations with weaker access. Allocation systems need explicit objectives, constraint review, and equity metrics. Decision makers should know whether the model maximizes throughput, clinical benefit, financial margin, or another target. Hidden objectives are dangerous because they transform institutional priorities into apparently neutral mathematics.

Clinical decision support is the point where these issues converge. Effective systems present relevant information at the moment of decision, communicate uncertainty, and make it easy for clinicians to accept, modify, or reject recommendations. Poor systems generate excessive alerts, obscure their reasoning, or interrupt care without offering a practical action. The literature favors collaborative intelligence over replacement (Shortliffe & Sepúlveda, 2018; Sutton *et al.*, 2020; Topol, 2019)^[34, 35, 38]. Clinicians remain responsible for context, patient preference, atypical presentations, and moral judgment. AI can extend attention and consistency, but it should not become an unchallengeable authority. External failures of proprietary sepsis tools and evidence of racial bias in population-health algorithms illustrate why local validation and subgroup analysis are essential (Obermeyer *et al.*, 2019)^[28].

Implementation Framework

A practical adoption pathway begins with problem selection, not technology procurement. Organizations should identify a measurable care or operational gap, map the existing workflow, and specify who will act on the output. The next stage is evidence review: data provenance, comparator choice, validation setting, calibration, subgroup performance, security testing, and reporting quality should be examined before purchase or development. Local silent testing should then evaluate the model without influencing care, revealing data-mapping errors and performance drift. A controlled pilot can assess usability, alert burden, response time, override patterns, and unintended effects. Expansion should occur only when predefined safety and value thresholds are met.

Governance should be multidisciplinary and continuous. A standing committee should include clinicians, nurses, data scientists, information-security staff, operations leaders, legal or compliance specialists, and patient representatives. Its responsibilities should cover intended use, version control, access rights, incident review, retraining, vendor changes, and retirement. Monitoring dashboards should track discrimination, calibration, missingness, subgroup outcomes, alert acceptance, downstream interventions, and adverse events. Performance should be reviewed after changes in clinical practice, population composition, coding, equipment, or software. Where evidence deteriorates, the system should be recalibrated, restricted, or withdrawn.

Economic evaluation also needs greater discipline. Studies frequently infer savings from prevented admissions, reduced reading time, or improved staffing, but implementation introduces integration, licensing, cybersecurity, training, maintenance, and governance costs. A credible business case should compare total cost with observed changes in outcomes and workload. Benefits may include avoided harm, faster diagnosis, reduced clinician burden, improved capacity, and resilience, not merely direct revenue. Hasan *et al.* (2025a)^[11] provide a useful cost-oriented framing, but prospective studies are needed to determine which savings persist outside modeled scenarios.

Equity and explainability should be treated as design requirements rather than final audits. Training data often reflect unequal access, underdiagnosis, and historical spending. If these proxies are used uncritically, AI can reproduce inequality at scale. Subgroup analysis should consider race, ethnicity, sex, age, disability, language, insurance status, geography, and relevant clinical factors, while recognizing that categories may be incomplete or socially constructed. Explainability should be matched to the user and decision. Clinicians may need influential variables, uncertainty, and comparable cases; patients may need a plain-language account of how information affected their care. Explanations should support scrutiny, not create false confidence.

Overall, the review suggests a hierarchy of readiness. Imaging triage, constrained risk prediction, anomaly detection, and demand forecasting are relatively mature when local validation is strong. Treatment recommendation, autonomous diagnosis, and high-stakes allocation require more prospective and comparative evidence. Generative systems that summarize records or produce recommendations add further risks of fabrication and source ambiguity.

The central policy implication is that healthcare AI should be regulated and managed according to intended use, potential harm, and degree of autonomy. Responsible adoption is

slower than installing software, but faster than recovering from a preventable clinical, ethical, or cybersecurity failure.

Table 3: Implementation risk-control matrix for healthcare AI

Risk	How it appears	Minimum control	Monitoring indicator
Distribution shift	Performance changes across time, sites, devices, or populations	Silent local testing, drift detection, recalibration	Calibration error and subgroup performance
Bias and inequity	Different error rates or resource access across groups	Representative data, fairness review, constraint testing	Outcome and error gaps by subgroup
Automation bias	Users defer to incorrect or overconfident outputs	Uncertainty display, override, independent review	Override rate and discordant-case outcomes
Alert fatigue	Excessive low-value notifications	Threshold tuning and workflow-centered design	Alert acceptance and response time
Privacy leakage	Sensitive data exposed through storage or model behavior	Minimization, encryption, federated or privacy-preserving methods	Access anomalies and leakage testing
Adversarial manipulation	Evasion, poisoning, model theft, or compromised devices	Red-team testing, signed models, secure updates, rollback	Security events and integrity checks
Uncertain value	Accuracy does not improve outcomes or reduce total cost	Prospective comparison and full economic evaluation	Patient outcomes, workload, and total cost
Vendor opacity	Undisclosed data, updates, or failure thresholds	Contractual transparency, audit access, version control	Change logs and independent validation

Conclusion

Artificial intelligence is reshaping healthcare through earlier prediction, faster image interpretation, stronger cyber defense, improved resource planning, and more informed clinical decisions. Across the 35 studies reviewed, the technology showed potential, but its value depended less on algorithmic novelty than on how well it was validated, integrated, governed, and monitored. Predictive models were most useful when linked to clear interventions. Imaging systems were strongest as second readers or prioritization tools. Cybersecurity applications required privacy and adversarial resilience throughout the model lifecycle. Resource-allocation tools needed explicit objectives and equity safeguards. Clinical decision support worked best when clinicians could understand uncertainty, retain authority, and act without excessive workflow burden. The evidence does not support replacing professional judgment with autonomous systems across complex care settings. Instead, it supports responsible collaboration in which AI expands human attention, consistency, and analytical capacity. Health systems should require prospective validation, subgroup assessment, cybersecurity-by-design, transparent reporting, economic evaluation, human override, and post-deployment surveillance. These safeguards are not obstacles to innovation; they are the conditions that make innovation clinically credible. The next phase of healthcare AI should therefore move beyond demonstrations of accuracy toward evidence of safer care, fairer access, resilient operations, and measurable patient benefit across real-world environments.

Limitations and Future Directions

This review has several limitations. The included studies were heterogeneous in population, task, data source, outcome definition, and validation design, preventing a meaningful

pooled effect estimate. The evidence base contained many retrospective and single-site studies, so reported performance may overstate real-world effectiveness. Publication bias may favor successful models, while negative implementations, withdrawn systems, and vendor evaluations are less likely to appear in journals. Some cybersecurity evidence came from adjacent financial or digital-infrastructure settings because healthcare-specific evaluations remain limited. The review included English-language publications and relied on publicly accessible scholarly records, which may have excluded relevant regional research. Rapid publication also means that new trials, regulations, and generative-AI applications may emerge after the search date.

Future research should prioritize prospective, multicenter, and pragmatic studies that compare AI-supported care with current practice using patient outcomes, workload, equity, and total cost. Models should be tested under temporal, geographic, demographic, and technical distribution shifts. Researchers should report calibration, uncertainty, subgroup performance, missing-data behavior, override patterns, and adverse events, not only accuracy. Cybersecurity studies should evaluate poisoning, evasion, model extraction, credential compromise, and recovery procedures in realistic clinical environments. Resource-allocation research should make objectives and fairness constraints explicit. Shared benchmark datasets, privacy-preserving collaboration, standardized incident reporting, and independent post-market surveillance would improve comparability. Future updates should also compare regulatory pathways and implementation outcomes across national health systems. Finally, patients and frontline professionals should participate in design and governance so that future systems solve meaningful problems, communicate limitations honestly, and remain accountable to the people affected.

Appendix A: Evidence Matrix of the 35 Included Studies

Study	Year	Primary domain	Task/data	Main contribution	Principal limitation
Miotto <i>et al.</i> (2016)	2016	Predictive analytics	EHR representation learning	Unsupervised patient representations predicted future disease	Single-system retrospective data
Gulshan <i>et al.</i> (2016)	2016	Medical imaging	Retinal photographs	High discrimination for referable diabetic retinopathy	Clinical impact not prospectively tested
Esteva <i>et al.</i> (2017)	2017	Medical imaging	Dermatology images	Deep networks achieved dermatologist-level classification	Image selection and spectrum constraints
Weng <i>et al.</i> (2017)	2017	Predictive analytics	Primary-care records	Improved cardiovascular-risk prediction	Retrospective validation
De Fauw <i>et al.</i> (2018)	2018	Medical imaging	Retinal OCT	Separated tissue mapping from referral recommendation	Site and device transfer require testing
Komorowski <i>et al.</i> (2018)	2018	Resource allocation	ICU sepsis treatment	Reinforcement learning derived treatment policies	Observational counterfactual evaluation
Rajkomar <i>et al.</i> (2018)	2018	Predictive analytics	Raw EHR data	Scalable prediction of mortality, readmission, and stay	Implementation effects not established
Shickel <i>et al.</i> (2018)	2018	Predictive analytics	Deep EHR review	Mapped major architectures and clinical applications	Heterogeneous benchmarks
Shortliffe & Sepúlveda (2018)	2018	Clinical decision support	Clinical informatics	Defined augmentation-focused decision support	Conceptual rather than comparative trial
Ardila <i>et al.</i> (2019)	2019	Medical imaging	Low-dose chest CT	End-to-end lung-cancer screening model	Retrospective screening cohorts
Finlayson <i>et al.</i> (2019)	2019	Cybersecurity	Adversarial medical ML	Demonstrated feasible manipulation of medical models	Attack realism varies by workflow
Liu <i>et al.</i> (2019)	2019	Medical imaging	Systematic review/meta-analysis	Compared deep learning with health professionals	High study heterogeneity and bias
Obermeyer <i>et al.</i> (2019)	2019	Clinical decision support	Population-health algorithm	Exposed racial bias from cost proxy	Single commercial use case
Tomašev <i>et al.</i> (2019)	2019	Predictive analytics	Longitudinal EHR	Continuous acute-kidney-injury prediction	Alert actionability and false positives
Topol (2019)	2019	Clinical decision support	Narrative synthesis	Articulated human-machine collaboration	Broad evidence rather than one intervention
Wiens <i>et al.</i> (2019)	2019	Clinical decision support	Responsible ML roadmap	Specified safety and implementation principles	Framework needs local operationalization
Kaissis <i>et al.</i> (2020)	2020	Cybersecurity	Privacy-preserving ML	Reviewed secure and federated approaches	Communication and attack trade-offs
McKinney <i>et al.</i> (2020)	2020	Medical imaging	Mammography	International breast-screening evaluation	Workflow and prevalence differences
Nagendran <i>et al.</i> (2020)	2020	Medical imaging	Systematic review	Identified weak design and reporting in clinician comparisons	Rapidly evolving evidence
Rieke <i>et al.</i> (2020)	2020	Cybersecurity	Federated learning	Outlined collaborative learning without raw-data pooling	Governance and poisoning remain concerns
Sutton <i>et al.</i> (2020)	2020	Clinical decision support	CDSS overview	Synthesized benefits, risks, and success strategies	Narrative evidence
El-Bouri <i>et al.</i> (2021)	2021	Resource allocation	Patient-flow review	Mapped forecasting and operational applications	Limited prospective implementation studies
Hasan <i>et al.</i> (2021)	2021	Predictive analytics	Cancer population data	Linked ML with incidence, mortality, and screening disparities	Ecological and public-data limitations
Hasan <i>et al.</i> (2022)	2022	Cybersecurity	Healthcare infrastructure	Proposed ML-enabled resilient cyber defense	Framework-level evidence
Rasel <i>et al.</i> (2022)	2022	Resource allocation	Healthcare supply chain	Integrated forecasting, visibility, and resilience strategies	Limited longitudinal case evidence
Tello <i>et al.</i> (2022)	2022	Resource allocation	Inpatient bed demand	Forecasted weekly capacity needs	Single-center data
Hasan <i>et al.</i> (2023)	2023	Cybersecurity	Fraud and risk analytics	Transferred AI risk controls to critical infrastructure	Financial-sector evidence is indirect
Elhaddad & Hamam (2024)	2024	Clinical decision support	AI-CDSS review	Synthesized technologies and adoption challenges	Narrative review design
Gómez-Cabello <i>et al.</i> (2024)	2024	Clinical decision support	Primary-care scoping review	Mapped current clinical implementations	Few outcome-focused deployments
Khosravi <i>et al.</i> (2024)	2024	Resource allocation	Review of reviews	Connected AI with service-delivery decisions	Dependent on prior-review quality
Milon <i>et al.</i> (2024)	2024	Cybersecurity	Cybercrime review	Integrated technical, legal, and sector impacts	Developing-economy focus
Hasan <i>et al.</i> (2025a)	2025	Predictive analytics	Cost and outcomes review	Connected prediction with readmissions, staffing, and cost	Some savings were modeled
Hasan <i>et al.</i> (2025b)	2025	Predictive analytics	Bacterial genomes	Proposed calibrated, lineage-aware AMR prediction	Clinical deployment remains prospective
Hasan <i>et al.</i> (2025c)	2025	Medical imaging	Brain MRI	Lightweight CNN with K-fold validation	Binary task and limited dataset
Hasan <i>et al.</i> (2025d)	2025	Cybersecurity	Predictive security review	Synthesized anomaly and fraud analytics	Adjacent financial domain

References

1. Ardila D, Kiraly AP, Bharadwaj S, Choi B, Reicher JJ, Peng L, *et al.* End-to-end lung cancer screening with three-dimensional deep learning on low-dose chest computed tomography. *Nat Med.* 2019;25(6):954–61. doi:10.1038/s41591-019-0447-x
2. Collins GS, Reitsma JB, Altman DG, Moons KGM. Transparent reporting of a multivariable prediction model for individual prognosis or diagnosis (TRIPOD): The TRIPOD statement. *Ann Intern Med.* 2015;162(1):55–63. doi:10.7326/M14-0697
3. Cruz Rivera S, Liu X, Chan AW, Denniston AK, Calvert MJ. Guidelines for clinical trial protocols for interventions involving artificial intelligence: The SPIRIT-AI extension. *Nat Med.* 2020;26:1351–63. doi:10.1038/s41591-020-1037-7
4. De Fauw J, Ledsam JR, Romera-Paredes B, Nikolov S, Tomasev N, Blackwell S, *et al.* Clinically applicable deep learning for diagnosis and referral in retinal disease. *Nat Med.* 2018;24(9):1342–50. doi:10.1038/s41591-018-0107-6
5. El-Bouri R, Taylor T, Youssef A, Zhu T, Clifton DA. Machine learning in patient flow: A review. *Prog Biomed Eng.* 2021;3(2):022002. doi:10.1088/2516-1091/abddc5
6. Elhaddad M, Hamam S. AI-driven clinical decision support systems: An ongoing pursuit of potential. *Cureus.* 2024;16(4):e57728. doi:10.7759/cureus.57728
7. Esteva A, Kuprel B, Novoa RA, Ko J, Swetter SM, Blau HM, *et al.* Dermatologist-level classification of skin cancer with deep neural networks. *Nature.* 2017;542(7639):115–8. doi:10.1038/nature21056
8. Finlayson SG, Bowers JD, Ito J, Zittrain JL, Beam AL, Kohane IS. Adversarial attacks on medical machine learning. *Science.* 2019;363(6433):1287–9. doi:10.1126/science.aaw4399
9. Gómez-Cabello CA, Borna S, Pressman S, Haider SA, Haider CR, Forte AJ. Artificial-intelligence-based clinical decision support systems in primary care: A scoping review of current clinical implementations. *Eur J Investig Health Psychol Educ.* 2024;14(3):685–98. doi:10.3390/ejihpe14030045
10. Gulshan V, Peng L, Coram M, Stumpe MC, Wu D, Narayanaswamy A, *et al.* Development and validation of a deep learning algorithm for detection of diabetic retinopathy in retinal fundus photographs. *JAMA.* 2016;316(22):2402–10. doi:10.1001/jama.2016.17216
11. Hasan MN, Arman M, Bhuyain MMH, Chowdhury F, Bathula MK. Predictive analytics in healthcare: Strategies for cost reduction and improved outcomes in the USA. *Int J Innov Res Sci Stud.* 2025a;8(8):142–50. doi:10.53894/ijirss.v8i8.10559
12. Hasan MN, Bhuyain MMH, Chowdhury F. AI model that predicts antibiotic resistance from bacterial genomes using open sequencing data. *Front Comput Sci Artif Intell.* 2025b;4(3):33–43. doi:10.32996/fcsai.2025.4.3.4
13. Hasan MN, Bhuyain MMH, Chowdhury F, Arman M. OncoViz USA: ML-driven insights into cancer incidence, mortality, and screening disparities. *J Med Health Stud.* 2021;2(1):53–62. doi:10.32996/jmhs.2021.2.1.6
14. Hasan MN, Miah MS, Ghose P, Jannat T, Bhuyain MMH, Chowdhury MSA, *et al.* A deep learning approach for brain tumor diagnosis: Combining an 8-layer CNN with rigorous K-fold validation. *Math Model Eng Probl.* 2025c;12(12):4387–96. doi:10.18280/mmep.121227
15. Hasan MN, Papel MSI, Rasel IH, Akter S, Aktar MK, Abedin MZ, *et al.* Enhancing financial information security through advanced predictive analytics: A PRISMA-based systematic review. *Edelweiss Appl Sci Technol.* 2025d;9(7):2222–45. doi:10.55214/2576-8484.v9i7.9142
16. Hasan MN, Rasel IH, Arman M, Ibrahim M, Jahan N. Strengthening U.S. financial and cybersecurity infrastructure with AI-driven fraud detection and risk analytics. *J Comput Anal Appl.* 2023;31(2):15–32.
17. Hasan MN, Rasel IH, Rahman M, Islam K, Arman M, Jahan N. Securing U.S. healthcare infrastructure with machine learning: Protecting patient data as a national security priority. *Int J Comput Exp Sci Eng.* 2022;8(3):85–93. doi:10.22399/ijcesen.3987
18. Kaissis GA, Makowski MR, Rückert D, Braren RF. Secure, privacy-preserving and federated machine learning in medical imaging. *Nat Mach Intell.* 2020;2(6):305–11. doi:10.1038/s42256-020-0186-1
19. Khosravi M, Zare Z, Mojtavaeian SM, Izadi R. Artificial intelligence and decision-making in healthcare: A thematic analysis of a systematic review of reviews. *Health Serv Res Manag Epidemiol.* 2024;11:23333928241234863. doi:10.1177/23333928241234863
20. Komorowski M, Celi LA, Badawi O, Gordon AC, Faisal AA. The artificial intelligence clinician learns optimal treatment strategies for sepsis in intensive care. *Nat Med.* 2018;24(11):1716–20. doi:10.1038/s41591-018-0213-5
21. Liu X, Faes L, Kale AU, Wagner SK, Fu DJ, Bruynseels A, *et al.* A comparison of deep learning performance against health-care professionals in detecting diseases from medical imaging: A systematic review and meta-analysis. *Lancet Digit Health.* 2019;1(6):e271–97. doi:10.1016/S2589-7500(19)30123-2
22. Liu X, Cruz Rivera S, Moher D, Calvert MJ, Denniston AK. Reporting guidelines for clinical trial reports for interventions involving artificial intelligence: The CONSORT-AI extension. *Nat Med.* 2020;26:1364–74. doi:10.1038/s41591-020-1034-x
23. McKinney SM, Sieniek M, Godbole V, Godwin J, Antropova N, Ashrafian H, *et al.* International evaluation of an AI system for breast cancer screening. *Nature.* 2020;577(7788):89–94. doi:10.1038/s41586-019-1799-6
24. Milon MNU, Ghose P, Chowdhury Pinky T, Nowshin Tabassum M, Nazmul Hasan M, Khatun M. An in-depth PRISMA-based review of cybercrime in a developing economy: Examining sector-wide impacts, legal frameworks, and emerging trends in the digital era. *Edelweiss Appl Sci Technol.* 2024;8(4):2072–93. doi:10.55214/25768484.v8i4.1583
25. Miotto R, Li L, Kidd BA, Dudley JT. Deep Patient: An unsupervised representation to predict the future of patients from the electronic health records. *Sci Rep.* 2016;6:26094. doi:10.1038/srep26094
26. Moons KGM, Wolff RF, Riley RD, Whiting PF, Westwood M, Collins GS, *et al.* PROBAST: A tool to assess risk of bias and applicability of prediction model studies. *Ann Intern Med.* 2019;170(1):51–8. doi:10.7326/M18-1376

27. Nagendran M, Chen Y, Lovejoy CA, Gordon AC, Komorowski M, Harvey H, *et al.* Artificial intelligence versus clinicians: Systematic review of design, reporting standards, and claims of deep learning studies. *BMJ*. 2020;368:m689. doi:10.1136/bmj.m689
28. Obermeyer Z, Powers B, Vogeli C, Mullainathan S. Dissecting racial bias in an algorithm used to manage the health of populations. *Science*. 2019;366(6464):447–53. doi:10.1126/science.aax2342
29. Page MJ, McKenzie JE, Bossuyt PM, Boutron I, Hoffmann TC, Mulrow CD, *et al.* The PRISMA 2020 statement: An updated guideline for reporting systematic reviews. *BMJ*. 2021;372:n71. doi:10.1136/bmj.n71
30. Rajkomar A, Oren E, Chen K, Dai AM, Hajaj N, Hardt M, *et al.* Scalable and accurate deep learning with electronic health records. *npj Digit Med*. 2018;1:18. doi:10.1038/s41746-018-0029-1
31. Rasel IH, Arman M, Hasan MN, Bhuyain MMH. Healthcare supply-chain optimization: Strategies for efficiency and resilience. *J Med Health Stud*. 2022;3(4):171–82. doi:10.32996/jmhs.2022.3.4.26
32. Rieke N, Hancox J, Li W, Milletari F, Roth HR, Albarqouni S, *et al.* The future of digital health with federated learning. *npj Digit Med*. 2020;3:119. doi:10.1038/s41746-020-00323-1
33. Shickel B, Tighe PJ, Bihorac A, Rashidi P. Deep EHR: A survey of recent advances in deep learning techniques for electronic health record analysis. *IEEE J Biomed Health Inform*. 2018;22(5):1589–604. doi:10.1109/JBHI.2017.2767063
34. Shortliffe EH, Sepúlveda MJ. Clinical decision support in the era of artificial intelligence. *JAMA*. 2018;320(21):2199–200. doi:10.1001/jama.2018.17163
35. Sutton RT, Pincock D, Baumgart DC, Sadowski DC, Fedorak RN, Kroeker KI. An overview of clinical decision support systems: Benefits, risks, and strategies for success. *npj Digit Med*. 2020;3:17. doi:10.1038/s41746-020-0221-y
36. Tello M, Reich ES, Puckey J, Maff R, Garcia-Arce A, Bhattacharya BS, *et al.* Machine learning based forecast for the prediction of inpatient bed demand. *BMC Med Inform Decis Mak*. 2022;22(1):55. doi:10.1186/s12911-022-01787-9
37. Tomašev N, Glorot X, Rae JW, Zielinski M, Askham H, Saraiva A, *et al.* A clinically applicable approach to continuous prediction of future acute kidney injury. *Nature*. 2019;572(7767):116–9. doi:10.1038/s41586-019-1390-1
38. Topol EJ. High-performance medicine: The convergence of human and artificial intelligence. *Nat Med*. 2019;25(1):44–56. doi:10.1038/s41591-018-0300-7
39. Vasey B, Nagendran M, Campbell B, Clifton DA, Collins GS, Denaxas S, *et al.* Reporting guideline for the early-stage clinical evaluation of decision support systems driven by artificial intelligence: DECIDE-AI. *Nat Med*. 2022;28(5):924–33. doi:10.1038/s41591-022-01772-9
40. Weng SF, Reys J, Kai J, Garibaldi JM, Qureshi N. Can machine-learning improve cardiovascular risk prediction using routine clinical data? *PLoS One*. 2017;12(4):e0174944. doi:10.1371/journal.pone.0174944
41. Whiting PF, Rutjes AWS, Westwood ME, Mallett S, Deeks JJ, Reitsma JB, *et al.* QUADAS-2: A revised tool for the quality assessment of diagnostic accuracy studies. *Ann Intern Med*. 2011;155(8):529–36. doi:10.7326/0003-4819-155-8-201110180-00009
42. Wiens J, Saria S, Sendak M, Ghassemi M, Liu VX, Doshi-Velez F, *et al.* Do no harm: A roadmap for responsible machine learning for health care. *Nat Med*. 2019;25(9):1337–40. doi:10.1038/s41591-019-0548-6

How to Cite This Article

Andrevna A, Juel D. Artificial intelligence in healthcare: a comprehensive systematic review of predictive analytics, medical imaging, cybersecurity, resource allocation, and clinical decision support. *Int J Artif Intell Eng Transform*. 2026;7(2):14-24. doi:10.54660/IJAET.2026.7.2.14-24.

Creative Commons (CC) License

This is an open access journal, and articles are distributed under the terms of the Creative Commons Attribution NonCommercial-ShareAlike 4.0 International (CC BYNC-SA 4.0) License, which allows others to remix, tweak, and build upon the work non-commercially, as long as appropriate credit is given and the new creations are licensed under the identical terms.