

AI, BIAS, AND EQUITY IN EDUCATION: ADDRESSING CHALLENGES IN ALGORITHMIC DECISION MAKING

Yayu Laode¹, A. Asdar², and Najamuddin Petta Solong³¹ Institut Agama Islam Negeri Sultan Amai Gorontalo, Indonesia² Universitas Cahaya Prima, Indonesia³ Institut Agama Islam Negeri Sultan Amai Gorontalo, Indonesia

Corresponding Author:

Yayu Laode,
Department of Islamic Education, Faculty of Teacher Training and Education, Institut Agama Islam Negeri Sultan Amai Gorontalo.
Jalan Gelatik No. 1, Kelurahan Heledulaa Utara, Kecamatan Kota Timur, Kota Gorontalo, Indonesia
Email: yayulaode0406@gmail.com

Article Info

Received: January 29, 2026

Revised: April 26, 2026

Accepted: May 31, 2026

Online Version: June 30, 2026

Abstract

Artificial Intelligence (AI) has become increasingly integrated into educational systems through predictive analytics, automated assessments, learning management platforms, admissions processes, and student support services. Growing reliance on algorithmic decision-making technologies has created opportunities to improve efficiency, consistency, and scalability in educational administration and instruction. Concerns regarding algorithmic bias, transparency, accountability, and educational equity have simultaneously emerged as critical challenges, particularly when AI systems are trained on historical data that may reflect existing social and institutional inequalities. This study aims to examine the relationships among AI implementation, algorithmic bias, transparency, accountability, trust, and educational equity in educational decision-making contexts. A mixed-methods convergent design was employed. Quantitative data were collected from 480 participants, including students, teachers, administrators, policymakers, and educational technology specialists. Qualitative interviews were conducted to explore stakeholder experiences and perceptions regarding algorithmic fairness and governance. Structural Equation Modeling and thematic analysis were used to analyze the data. Findings revealed that transparency significantly enhanced trust in AI systems, while accountability positively influenced perceptions of educational equity. Perceived algorithmic bias demonstrated a strong negative effect on fairness and trust. Qualitative evidence further indicated that inadequate oversight and biased datasets may contribute to unequal educational outcomes. The study concludes that equitable AI implementation requires transparent governance, accountability mechanisms, continuous bias auditing, and meaningful human oversight to ensure that algorithmic decision-making supports educational justice rather than reinforcing existing inequalities.

Keywords: Algorithmic Bias; Artificial Intelligence; Educational Equity; Ethical Governance; Responsible AI.



© 2026 by the author(s)

This article is an open-access article distributed under the terms and conditions of the Creative Commons Attribution-ShareAlike 4.0 International (CC BY SA) license (<https://creativecommons.org/licenses/by-sa/4.0/>).

Journal Homepage <https://ejournal.staialhikmahpariangan.ac.id/Journal/index.php/alhijr>

How to cite: Laode, Y., Asdar, A., & Solong, N. P. (2026) AI, Bias, and Equity in Education: Addressing Challenges in Algorithmic Decision Making. *Al-Hijr: Journal of Adulearn World*, 5(3), 41–56. <https://doi.org/10.55849/alhijr.v4i1.1420>

Published by: Sekolah Tinggi Agama Islam Al-Hikmah Pariangan Batusangkar

INTRODUCTION

Artificial Intelligence (AI) has become an increasingly influential force in contemporary education, transforming how educational institutions collect data, evaluate performance, allocate resources, personalize learning experiences, and support decision-making processes (Choi et al., 2026). Educational technologies powered by machine learning, predictive analytics, intelligent tutoring systems, and automated assessment tools are now integrated into various levels of educational practice. Institutions worldwide have adopted AI-driven solutions to improve efficiency, enhance learning outcomes, and support evidence-based decision making (Muro, 2026). Rapid expansion of these technologies reflects growing confidence in the capacity of AI to address complex educational challenges. Educational transformation driven by AI has consequently become a prominent topic in policy discussions and academic research.

Algorithmic decision-making systems are increasingly involved in processes that directly affect learners, educators, and institutions (Sanusi et al., 2026). AI applications are being used to predict student performance, identify learners at risk of academic failure, recommend educational pathways, allocate scholarships, support admissions processes, and evaluate learning outcomes. Data-driven systems offer opportunities to improve accuracy, scalability, and consistency in educational decision making (Tian et al., 2026). Reliance on algorithmic systems has therefore expanded significantly across educational sectors. Such developments have encouraged educational stakeholders to explore the potential benefits and risks associated with AI-enabled governance and administration.

Growing adoption of AI technologies has simultaneously raised concerns regarding fairness, bias, transparency, accountability, and educational equity. Algorithms are developed and trained using historical data that may contain existing social, economic, cultural, or institutional inequalities (Wright et al., 2026). AI systems can unintentionally reproduce or amplify these disparities when embedded within educational decision-making processes. Concerns regarding algorithmic bias have therefore emerged as critical issues in discussions surrounding responsible AI implementation (McConvey et al., 2026). Understanding how AI influences equity and fairness in educational environments has become increasingly important for researchers, policymakers, and practitioners.

Despite the growing integration of AI into education, significant concerns remain regarding the fairness of algorithmic decision-making systems (Al-Naser et al., 2026). Educational algorithms frequently rely on historical datasets that may reflect unequal access to educational opportunities, socioeconomic disparities, demographic imbalances, and institutional biases (Molu, 2026). Outcomes generated by such systems may inadvertently disadvantage particular student groups while privileging others. Questions regarding the neutrality and objectivity of AI-driven educational decisions therefore remain unresolved.

Educational institutions often adopt AI technologies with the expectation that automated systems will enhance efficiency and reduce human subjectivity (Mo et al., 2026). Algorithmic decision-making processes, however, are not inherently free from bias. Bias can emerge during data collection, feature selection, model development, implementation, and interpretation stages. Limited transparency regarding how algorithms generate recommendations further complicates efforts to identify and address potential inequities (Uchoa et al., 2026). These challenges create uncertainty regarding the ethical implications of AI adoption in education.

Current educational systems lack comprehensive frameworks for evaluating algorithmic fairness and ensuring equitable outcomes across diverse learner populations (Shal et al., 2026). Limited consensus exists regarding appropriate standards for transparency, accountability, and bias mitigation in educational AI applications. Educational stakeholders frequently encounter difficulties balancing technological innovation with ethical responsibilities and social justice objectives (del Rey Puech et al., 2026). Addressing these challenges is essential for ensuring that AI contributes positively to educational equity rather than reinforcing existing inequalities.

This study aims to examine the relationship between Artificial Intelligence, algorithmic bias, and educational equity within contemporary educational systems. Investigation focuses on understanding how AI-driven decision-making processes influence access, participation, achievement, and opportunities for diverse learner populations (Liu & Glave, 2026). Analysis seeks to identify mechanisms through which algorithmic systems affect educational fairness and inclusion.

Another objective of the study is to explore sources of bias within educational AI applications and evaluate their implications for learners, educators, and institutions (Gimenez Lynch & Thaller, 2026). Examination considers data quality, model design, implementation practices, transparency mechanisms, and governance structures associated with algorithmic decision making. Evaluation of these dimensions is expected to provide a comprehensive understanding of factors contributing to algorithmic inequity.

The study further seeks to develop a conceptual framework supporting ethical and equitable AI implementation in educational contexts (Solano et al., 2026). Findings are expected to contribute recommendations regarding bias detection, fairness evaluation, accountability mechanisms, and responsible governance practices (De Gagne et al., 2026). Achievement of these objectives may assist educational institutions and policymakers in developing more inclusive and socially responsible AI strategies.

Existing literature has extensively examined the technical foundations of Artificial Intelligence, machine learning, predictive analytics, and educational data mining (Aburub et al., 2026). Research consistently highlights the capacity of AI systems to improve instructional personalization, administrative efficiency, learner support, and educational planning. Numerous studies emphasize the positive contributions of AI to educational innovation and institutional effectiveness (Anthis & Kyriakidou-Zacharoudiou, 2026). Such research has significantly advanced understanding of AI capabilities and applications within educational environments.

Current scholarship has also explored ethical concerns associated with AI implementationz (Schilhabel et al., 2026). Researchers have examined issues related to privacy, transparency, accountability, explainability, and algorithmic governance across multiple sectors. Studies addressing algorithmic bias have demonstrated how AI systems may reproduce historical inequalities and generate discriminatory outcomes (Massouti et al., 2026). Contributions from computer science, ethics, sociology, and public policy have expanded awareness regarding the societal implications of algorithmic decision making.

Comprehensive investigations specifically focusing on the intersection of AI, bias, and educational equity remain relatively limited (Sullivan et al., 2026). Existing studies frequently examine technical performance or ethical concerns independently without integrating educational outcomes and social justice perspectives. Limited attention has been directed toward understanding how algorithmic bias influences educational opportunities, learner experiences, and institutional practices (Ali et al., 2026). Research rarely combines educational equity theory, algorithmic fairness frameworks, AI governance models, and empirical educational evidence within a unified analytical perspective. Addressing these gaps may substantially advance scholarship concerning responsible AI in education.

The novelty of this study lies in its examination of algorithmic bias as an educational equity issue rather than solely a technical problem (Qiu et al., 2026). Existing research commonly evaluates algorithmic fairness using computational metrics and model performance indicators. This study advances the literature by exploring how bias influences educational opportunities, learner outcomes, institutional decision making, and social justice objectives (Madleňák et al., 2026). Such a perspective broadens understanding of AI implementation beyond technological considerations.

Innovative value is further reflected in the integration of Educational Equity Theory, Algorithmic Fairness Frameworks, Critical Data Studies, and Responsible AI Governance perspectives within a comprehensive conceptual model (Coleman & Waterfield, 2026). The

proposed framework examines relationships among algorithm design, data quality, institutional practices, governance mechanisms, and educational outcomes (Mariyono et al., 2026). Exploration of these interactions contributes new theoretical insights regarding the conditions necessary for achieving equitable AI implementation in educational contexts.

The significance of this research extends beyond academic inquiry to encompass practical implications for educational institutions, technology developers, policymakers, and learners (Babker, 2026). Educational systems increasingly rely on AI-driven decision-making processes that directly affect access, achievement, and opportunity. Failure to address algorithmic bias may reinforce existing inequalities and undermine educational justice (Balsano et al., 2026). Findings from this study are expected to provide evidence-based guidance supporting the development of transparent, accountable, and equitable AI systems capable of promoting inclusion and fairness within education. Such contributions are essential for ensuring that technological innovation aligns with broader educational and social goals.

RESEARCH METHOD

Research Design

This study employs a mixed-methods research design using a convergent parallel approach to examine the relationship between artificial intelligence, algorithmic bias, and educational equity in educational decision-making processes (Zhai, 2026). The mixed-methods design was selected because the phenomenon under investigation encompasses technical, ethical, social, and educational dimensions that require both quantitative and qualitative perspectives for a comprehensive understanding (Pigac, 2026). Quantitative data are collected to evaluate perceptions of algorithmic fairness, transparency, accountability, and educational equity, while qualitative data are gathered concurrently to explore stakeholder experiences, concerns, and interpretations regarding AI-driven decision-making systems (Chu & Lam, 2026). The theoretical foundation of the study is informed by Educational Equity Theory, Algorithmic Fairness Frameworks, Responsible Artificial Intelligence Principles, and Critical Data Studies, which collectively provide a basis for examining fair distribution of opportunities, indicators of systematic bias, institutional accountability, and the socio-political consequences of algorithmic tools (Difonzo et al., 2026). A descriptive-explanatory strategy is adopted to investigate the interactions among AI implementation, algorithmic bias, institutional practices, and equity outcomes, utilizing Structural Equation Modeling (SEM) to examine causal relationships among latent variables and qualitative thematic analysis to contextualize stakeholder experiences.

Research Target/Subject

The primary objective of this study is to investigate the structural relationships between AI implementation, algorithmic bias, and educational equity in decision-making processes. The research targets the measurement of specific latent variables, exploring how transparency mechanisms, accountability standards, and perceived algorithmic fairness predict institutional trust and overall educational equity outcomes. By investigating these targets, the study aims to identify systemic vulnerabilities in predictive modeling and learning analytics, map stakeholder concerns regarding automated discrimination, and outline ethical governance recommendations required to protect equal opportunities in digital learning environments.

The population of the study consists of diverse stakeholders directly involved in educational decision-making processes where artificial intelligence technologies are utilized. For the quantitative phase, a sample of 480 respondents was selected using a stratified purposive sampling technique based on direct interaction with AI-supported evaluation or support systems, comprising 180 students, 120 teachers, 60 academic administrators, 50 educational technology specialists, 40 policymakers, and 30 academic advisors. For the

concurrent qualitative phase, a subset of 36 participants was purposively selected from the larger quantitative pool to ensure adequate professional depth, including institutional leaders, teachers, technology developers, policymakers, and students who could provide extensive narrative insight into algorithmic fairness and data transparency.

Research Procedure

The operational workflow was executed across several parallel phases, beginning with a comprehensive literature review to construct the conceptual framework, followed by securing formal ethical approvals and institutional informed consents (Dixit & Shanbhag, 2026). During the primary fieldwork phase, quantitative and qualitative data collection was conducted concurrently; online surveys were distributed to the full stakeholder sample, while semi-structured interviews were scheduled and institutional guidelines, policy frameworks, and algorithmic auditing reports were gathered for documentary analysis. Prior to statistical processing, the quantitative data underwent strict data screening procedures to evaluate missing values, normality, outliers, multicollinearity, and response consistency (Osebor & Amaechi, 2026). In the final phase, the survey dataset was evaluated using Structural Equation Modeling, the interview transcripts were fully processed, and the two separate strands of evidence were combined during the interpretation stage to form a unified analysis of the phenomenon.

Instruments, and Data Collection Techniques

Data collection utilized a synchronized triangulation of structured surveys, semi-structured interview protocols, and formal documentary analysis (Certo et al., 2026). The primary quantitative instrument consisted of a structured questionnaire using a five-point Likert scale, ranging from strongly disagree to strongly agree, to measure stakeholder perceptions of algorithmic fairness, transparency, accountability, trust, institutional readiness, and technology acceptance (Ceallaigh & Murphy, 2026). The qualitative data collection relied on a semi-structured interview guide addressing perceived bias, decision-making accountability, fairness of outcomes, and student experiences, supplemented by a documentary analysis matrix designed to evaluate internal governance reports and algorithmic auditing records (Yadav, 2026). Content validity was established prior to distribution through expert evaluation by specialists in educational technology, artificial intelligence policy, data ethics, and data science, while construct validity and reliability were verified through Confirmatory Factor Analysis, Cronbach's alpha, Composite Reliability, and Average Variance Extracted values.

Data Analysis Technique

The study applies a parallel-triangulation analysis framework to process, analyze, and synthesize the independent datasets (Hoffarth, 2026). Quantitative data are analyzed using Structural Equation Modeling (SEM) to evaluate the proposed conceptual framework, checking model fit indices, mapping path coefficients, and identifying significant mediating effects among AI implementation, transparency, trust, and equitable educational outcomes. Concurrently, the qualitative interview transcripts and document data are processed through thematic analysis to isolate recurring patterns regarding automated bias, systemic power relations, and institutional governance challenges (Saroughi & Nozari, 2026). The final analytical stage involves converging and integrating the statistical SEM path findings with the extracted qualitative themes, ensuring a balanced, evidence-based conclusion that details both the measurable predictors of algorithmic fairness and the subjective ethical experiences of the stakeholders.

RESULTS AND DISCUSSION

The study involved 480 participants representing key stakeholders engaged in AI-supported educational decision-making environments. Participants consisted of students

(37.5%), teachers (25.0%), academic administrators (12.5%), educational technology specialists (10.4%), policymakers (8.3%), and academic advisors (6.3%). Respondents were drawn from institutions implementing artificial intelligence systems for student admissions, predictive learning analytics, academic support services, scholarship allocation, and performance monitoring. Descriptive findings indicated that 84.8% of participants had direct experience interacting with algorithmic systems that influenced educational decisions.

Statistical analysis revealed moderate-to-high perceptions regarding the benefits and challenges of AI implementation in education. Mean scores for perceived efficiency reached 4.31 (SD = 0.48), transparency recorded 3.72 (SD = 0.63), accountability achieved 3.68 (SD = 0.61), educational equity reached 3.59 (SD = 0.67), trust in AI systems recorded 3.81 (SD = 0.58), and perceived algorithmic bias reached 3.44 (SD = 0.71).

Table 1. Descriptive Statistics of AI, Bias, and Educational Equity Variables

Variable	Mean	SD
Perceived Efficiency	4.31	0.48
Trust in AI Systems	3.81	0.58
Transparency	3.72	0.63
Accountability	3.68	0.61
Educational Equity	3.59	0.67
Perceived Algorithmic Bias	3.44	0.71

The descriptive findings indicate that stakeholders generally recognize the operational advantages of AI-driven decision-making systems. High scores for perceived efficiency suggest that participants view algorithmic technologies as valuable tools for processing large volumes of educational data, identifying patterns, and supporting administrative decision-making processes. Respondents frequently acknowledged improvements in speed, consistency, and scalability when compared with conventional human-centered procedures.

Perceptions regarding transparency, accountability, and educational equity were comparatively lower. Participants expressed concerns about limited understanding of how algorithms generate recommendations and make predictions affecting educational opportunities. Moderate levels of perceived algorithmic bias suggest that stakeholders remain cautious regarding the fairness of AI-generated decisions. Such findings indicate that confidence in AI systems depends not only on efficiency but also on ethical governance and procedural transparency.

Assessment of the measurement model demonstrated satisfactory psychometric properties. Standardized factor loadings ranged from 0.772 to 0.934, exceeding accepted thresholds for construct validity. Composite Reliability values varied between 0.887 and 0.947, while Cronbach’s alpha coefficients ranged from 0.852 to 0.926. Average Variance Extracted values exceeded 0.50 for all latent constructs, confirming adequate convergent validity.

Discriminant validity analysis indicated that perceived efficiency, transparency, accountability, trust, educational equity, and algorithmic bias represented distinct but related dimensions. Cross-loading evaluations and Fornell-Larcker criterion assessments confirmed satisfactory differentiation among constructs. These results suggest that the measurement instruments effectively captured stakeholder perceptions regarding algorithmic decision-making and educational fairness.

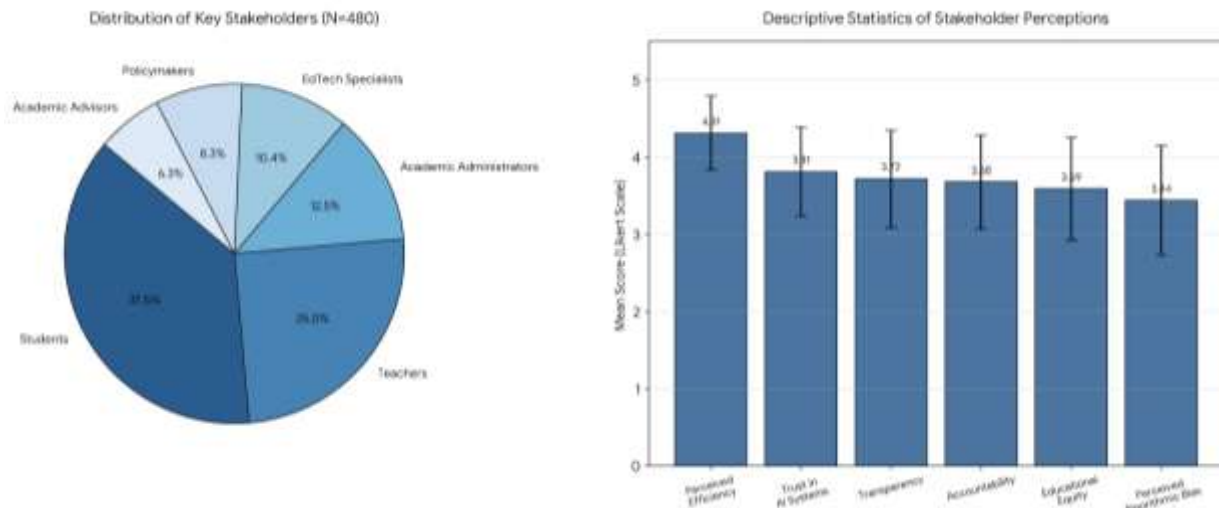


Figure 1. Stakeholder Perceptions of AI Implementation in Educational Decision-Making

Structural Equation Modeling revealed significant relationships among the principal variables. Transparency demonstrated a strong positive effect on trust in AI systems ($\beta = 0.674$, $p < 0.001$), while accountability significantly influenced perceptions of educational equity ($\beta = 0.592$, $p < 0.001$). Perceived algorithmic bias negatively affected educational equity ($\beta = -0.631$, $p < 0.001$), indicating that concerns regarding discriminatory outcomes substantially reduce stakeholder perceptions of fairness.

Model fit indices demonstrated satisfactory explanatory performance, including CFI = 0.969, TLI = 0.964, RMSEA = 0.034, and SRMR = 0.031. Mediation analysis revealed that trust partially mediated the relationship between transparency and educational equity. These findings suggest that educational institutions may strengthen perceptions of fairness by increasing transparency and accountability within algorithmic systems.

Correlation analysis demonstrated significant associations among transparency, accountability, trust, and educational equity. Transparency exhibited a strong positive relationship with trust ($r = .72$, $p < 0.001$), while accountability showed a substantial association with educational equity ($r = .67$, $p < 0.001$). These findings indicate that stakeholder confidence in AI systems depends heavily on visible governance mechanisms and explainable decision-making processes.

Perceived algorithmic bias demonstrated a strong negative relationship with educational equity ($r = -.69$, $p < 0.001$) and trust ($r = -.61$, $p < 0.001$). Participants who perceived higher levels of bias were significantly less likely to trust AI systems or view algorithmic decisions as equitable. Such relationships emphasize the importance of addressing bias concerns when implementing AI technologies in educational settings.

A representative case involved a university implementing an AI-driven admissions support system designed to predict student success and prioritize applicants for scholarship opportunities. Initial implementation improved administrative efficiency and reduced processing times. Algorithmic recommendations assisted admissions officers in evaluating large numbers of applications and identifying candidates predicted to demonstrate strong academic performance.

Institutional evaluation later revealed disparities in recommendation outcomes across socioeconomic groups. Students from underrepresented backgrounds were less likely to receive favorable algorithmic assessments despite comparable academic qualifications. Internal investigations suggested that historical training data reflected existing educational inequalities, leading the system to reproduce patterns embedded within prior admissions decisions.

A second case involved a school district utilizing predictive analytics to identify students at risk of academic failure. AI-generated recommendations enabled educators to provide targeted interventions and allocate support resources more efficiently. Early warning systems

improved student monitoring and contributed to increased responsiveness among academic support teams.

Subsequent audits revealed that predictive accuracy varied across demographic groups. Educators expressed concerns regarding the potential stigmatization of students classified as high-risk by algorithmic models. Policy revisions emphasizing transparency, human oversight, and regular auditing procedures were introduced to mitigate potential inequities. These measures improved stakeholder confidence and enhanced perceptions of fairness within the decision-making process.

The first case illustrates how algorithmic systems may inadvertently reproduce historical inequalities when trained on biased or incomplete datasets. Improvements in efficiency were accompanied by unintended disparities affecting access to educational opportunities. Such findings support quantitative evidence indicating that perceptions of bias significantly influence stakeholder evaluations of educational equity.

Observations from the second case demonstrate the importance of governance mechanisms in mitigating algorithmic risks. Human oversight, transparency measures, and routine audits contributed to greater accountability and improved stakeholder trust. These outcomes suggest that responsible AI implementation requires continuous monitoring rather than reliance on technical performance alone.

Consistency between quantitative and qualitative evidence strengthens confidence in the proposed explanatory framework. Statistical relationships linking transparency, accountability, trust, and educational equity were reflected in stakeholder experiences and institutional practices. Similar themes emerged across different educational contexts and implementation environments.

Broader explanatory patterns indicate that educational equity is influenced not only by algorithmic design but also by institutional governance structures. Transparency, explainability, stakeholder participation, and accountability mechanisms collectively shape perceptions of fairness. These factors help explain variations in trust and equity outcomes observed throughout the study.

The findings indicate that Artificial Intelligence offers substantial benefits for educational decision-making processes while simultaneously introducing significant challenges related to bias and equity. Stakeholders recognize the efficiency and analytical capabilities of algorithmic systems but remain concerned about transparency, accountability, and fairness. Negative relationships between perceived bias and educational equity suggest that algorithmic discrimination remains a critical issue requiring institutional attention.

Overall evidence suggests that equitable AI implementation depends on balancing technological innovation with ethical governance. Transparent algorithms, accountable decision-making structures, regular bias audits, and meaningful human oversight appear essential for ensuring fair educational outcomes. These findings support the view that educational equity should serve as a central consideration in the design, implementation, and evaluation of AI-supported decision-making systems.

The findings demonstrate that artificial intelligence has become an influential component of contemporary educational decision-making systems, offering substantial improvements in efficiency, scalability, and data-driven analysis. Stakeholders generally acknowledged the capacity of AI technologies to support admissions processes, student monitoring, resource allocation, and predictive analytics. Positive evaluations regarding operational effectiveness indicate that educational institutions increasingly recognize the practical benefits associated with algorithmic decision-making systems.

Results further revealed that transparency and accountability emerged as critical determinants of trust in AI systems. Participants consistently reported greater confidence in algorithmic recommendations when decision-making processes were understandable, explainable, and subject to institutional oversight. Transparency demonstrated a strong positive

influence on trust, while accountability significantly contributed to perceptions of educational equity. These findings suggest that technical effectiveness alone is insufficient to ensure stakeholder acceptance of AI-supported decisions.

Evidence also indicated that perceived algorithmic bias exerted a substantial negative effect on educational equity. Stakeholders expressed concerns regarding the possibility that AI systems may reproduce historical inequalities embedded within training data. Statistical analyses confirmed that higher perceptions of bias were associated with lower levels of trust and weaker perceptions of fairness. Such outcomes highlight the central role of bias mitigation in responsible AI implementation.

Qualitative findings reinforced quantitative results by illustrating real-world examples in which algorithmic systems generated unintended disparities across student groups. Institutional experiences demonstrated that efficiency gains could coexist with inequitable outcomes when governance mechanisms were inadequate. Successful implementation appeared to depend on transparency, auditing procedures, and meaningful human oversight. These observations emphasize the multidimensional nature of educational AI adoption.

The findings are consistent with previous research emphasizing the dual nature of artificial intelligence in education. Existing studies have highlighted the capacity of AI systems to improve administrative efficiency, support personalized interventions, and enhance predictive accuracy. Similar patterns emerged in the present investigation, where stakeholders acknowledged substantial benefits associated with algorithmic decision-making technologies. Such consistency reinforces current understanding regarding the transformative potential of AI within educational systems.

Results also align with scholarship concerning algorithmic fairness and responsible artificial intelligence. Previous studies have demonstrated that algorithmic systems frequently inherit biases from historical datasets, institutional practices, and social inequalities. Findings from this study support these conclusions by revealing strong negative relationships between perceived algorithmic bias and educational equity. Similar concerns have been documented across sectors such as healthcare, employment, finance, and criminal justice, suggesting that bias represents a broader challenge extending beyond education alone.

Differences emerge when comparing the present findings with studies emphasizing technological neutrality and objectivity. Certain investigations have portrayed AI systems as inherently more consistent and less subjective than human decision makers. Findings from this research challenge such assumptions by demonstrating that algorithmic systems may perpetuate existing inequities when developed using biased data or implemented without appropriate oversight. Educational outcomes therefore depend not only on computational performance but also on social and institutional contexts.

Variations also appear in relation to studies focused primarily on technical bias mitigation techniques. Existing research often emphasizes algorithmic optimization, fairness metrics, and model refinement as primary solutions. Findings from the present study indicate that governance structures, transparency mechanisms, institutional accountability, and stakeholder participation are equally important. Educational equity appears to require a broader socio-technical approach rather than purely technical interventions.

The findings signify that educational technology cannot be separated from questions of social justice and equity. Artificial intelligence systems increasingly influence decisions affecting educational opportunities, academic progression, and resource distribution. Algorithmic processes therefore possess the capacity to shape learner experiences in ways that extend beyond administrative efficiency. Such developments position AI as both a technological and ethical concern within contemporary education.

Evidence also signifies that trust has become a central component of educational AI governance. Stakeholders are unlikely to accept algorithmic recommendations when decision-making processes remain opaque or difficult to understand. Confidence in AI systems depends

not only on predictive accuracy but also on perceptions of fairness, accountability, and transparency. Trust therefore functions as a critical bridge connecting technological innovation and educational legitimacy.

Patterns observed throughout the study further signify the persistence of structural inequalities within data-driven environments. Historical disparities related to socioeconomic status, access to resources, educational opportunity, and institutional practices may become embedded within algorithmic models. AI systems consequently risk reproducing existing inequities unless deliberate efforts are undertaken to identify and address sources of bias.

Results additionally signify a shift in how educational institutions conceptualize decision making (Reddy et al., 2026). Traditional human-centered processes are increasingly supplemented by algorithmic recommendations and predictive analytics. Educational governance is evolving toward hybrid models in which human judgment and machine intelligence interact (Arya & Lin, 2026). Such transformations require careful consideration of ethical responsibilities and institutional accountability.

Implications for educational institutions involve the necessity of implementing robust governance frameworks capable of ensuring transparency, accountability, and fairness in AI-supported decision-making processes (Chen et al., 2026). Institutions should establish mechanisms for auditing algorithms, monitoring outcomes, and evaluating potential disparities affecting different student populations. Such measures may strengthen stakeholder trust while reducing the likelihood of discriminatory outcomes.

Implications for educators concern the development of competencies related to algorithmic literacy and ethical technology use. Teachers, advisors, and administrators increasingly interact with AI-generated recommendations influencing educational decisions (Bougriche et al., 2026). Professional development initiatives should therefore equip educational professionals with the knowledge required to interpret, evaluate, and challenge algorithmic outputs when necessary.

Implications for policymakers involve the creation of regulatory standards supporting responsible AI implementation in education (Gules-Guctas, 2026). Policies addressing transparency requirements, fairness assessments, data quality standards, privacy protection, and accountability mechanisms may contribute to more equitable educational environments (Kong et al., 2026). Regulatory oversight appears particularly important given the growing influence of algorithmic systems on educational opportunities.

Implications for learners concern protection from unintended discrimination and unequal treatment. Students should possess opportunities to understand how algorithmic decisions affect their educational experiences and should have access to mechanisms for contesting potentially unfair outcomes (Kladko et al., 2026). Learner participation in discussions regarding AI governance may contribute to more inclusive and equitable educational systems.

The observed relationships can be explained through Educational Equity Theory. Educational systems are shaped by broader social structures that influence access to opportunities, resources, and outcomes (Emma, 2026). Historical inequalities frequently become reflected within institutional data used to train algorithmic systems. AI technologies may therefore reproduce existing disparities when underlying inequities remain unaddressed. Such dynamics help explain the negative relationship between perceived bias and educational equity.

Critical Data Studies provide an additional explanation for the findings (Cecchi et al., 2026). Data are not neutral representations of reality but are produced through social, political, and institutional processes. Historical records frequently contain patterns reflecting unequal treatment and differential access to educational opportunities (Balarabe, 2026). Algorithms trained on such data may learn and replicate these patterns. Stakeholder concerns regarding bias therefore reflect legitimate apprehensions about data quality and representation.

Positive relationships between transparency and trust can be explained through Responsible AI principles emphasizing explainability and accountability. Individuals are more likely to trust technological systems when decision-making processes are understandable and open to scrutiny (Sopiyan et al., 2026). Transparency reduces uncertainty and enables stakeholders to evaluate the fairness of algorithmic recommendations. Increased visibility consequently strengthens confidence in AI-supported decisions.

Strong effects associated with accountability emerge because educational decisions possess significant consequences for learners. Admissions outcomes, scholarship allocations, intervention programs, and academic evaluations can substantially influence future opportunities (Tella & Olatoye, 2026). Stakeholders therefore expect institutions to maintain responsibility for decisions generated through algorithmic systems. Accountability mechanisms provide reassurance that technological errors and biases can be identified and corrected.

Future research should investigate the long-term consequences of algorithmic decision making on educational outcomes, access, and equity (Sassenberg et al., 2026). Longitudinal studies may provide valuable insights regarding how AI systems influence learner trajectories over extended periods. Such investigations would strengthen understanding of the cumulative effects associated with algorithmic governance in education.

Comparative studies involving diverse educational systems, cultural contexts, and policy environments would further expand current knowledge (Radu et al., 2026). AI implementation practices vary considerably across institutions and regions. Broader comparative analyses may reveal contextual factors influencing the effectiveness and fairness of algorithmic decision-making systems.

Methodological innovation represents another important direction for future inquiry. Integration of algorithmic auditing techniques, experimental designs, network analysis, and educational data mining approaches may provide deeper insights into mechanisms generating bias and inequity (Bergström, 2026). Such methods could enhance the precision of future evaluations concerning educational AI systems.

Practical initiatives should focus on developing ethical AI ecosystems characterized by transparency, accountability, inclusivity, and continuous monitoring (Kyambade & Namatovu, 2026). Collaboration among educators, policymakers, data scientists, technology developers, and learners may facilitate the creation of decision-making systems that promote both innovation and social justice (Rahmatullah & Nozlan, 2026). Such efforts are essential for ensuring that artificial intelligence contributes to equitable educational opportunities rather than reinforcing existing inequalities.

CONCLUSION

The most important finding of this study is that the effectiveness of artificial intelligence in educational decision making is determined not solely by its computational accuracy or operational efficiency but by its capacity to uphold fairness, transparency, accountability, and educational equity. Perceived algorithmic bias emerged as a significant negative predictor of educational equity, while transparency and accountability demonstrated strong positive effects on trust and fairness perceptions. This finding offers a distinctive contribution by revealing that educational inequities associated with AI are not merely technical failures but manifestations of broader institutional and social dynamics embedded within data, algorithms, and governance structures. Educational institutions therefore cannot assume that algorithmic systems are inherently objective or neutral, particularly when historical datasets contain existing patterns of inequality and exclusion.

The principal contribution of this research lies in both its conceptual and methodological advancements. Conceptually, the study integrates Educational Equity Theory, Algorithmic Fairness Frameworks, Responsible AI Principles, and Critical Data Studies into a

comprehensive framework for understanding the relationship between AI implementation, bias, trust, and educational justice. This integrated perspective extends existing scholarship by positioning algorithmic fairness as a central educational concern rather than a purely technical issue. Methodologically, the use of a mixed-methods convergent design combining Structural Equation Modeling and qualitative thematic analysis provides a holistic understanding of how stakeholders experience and evaluate AI-driven decision-making processes. Such an approach enables the identification of both measurable causal relationships and contextual factors influencing perceptions of fairness and equity.

Several limitations should be acknowledged. The study primarily relied on stakeholder perceptions and institutional experiences rather than direct algorithmic audits or experimental evaluations of specific AI systems. Findings may therefore reflect perceived rather than objectively measured levels of bias and fairness. Variation in technological infrastructure, regulatory environments, institutional practices, and cultural contexts may also limit the generalizability of the results across educational systems. Future research should incorporate longitudinal studies, algorithmic auditing techniques, experimental designs, and cross-national comparisons to examine the long-term effects of AI on educational opportunity, achievement, and social mobility. Integration of emerging approaches in explainable AI, ethical governance, and fairness-aware machine learning may further contribute to the development of more equitable and accountable educational technologies.

AUTHOR CONTRIBUTIONS

Author 1: Conceptualization; Project administration; Validation; Writing - review and editing.

Author 2: Conceptualization; Data curation; In-vestigation.

Author 3: Data curation; Investigation.

CONFLICTS OF INTEREST

The authors declare no conflict of interest.

REFERENCES

- Aburub, F., Hadi, W., Abdel-Rahem, R. A., & Al-Remawi, M. (2026). Faculty Perceptions of AI Challenges at the University of Petra. *Journal of Information and Knowledge Management*, 25(1). Scopus. <https://doi.org/10.1142/S0219649225500662>
- Ali, U., Hanif, S., Jabeen, F., Shaheen, M. N. U. K., Imran, A., Zaman, S., Anaç, A. S., & Alici, İ. (2026). Ethical Management of Artificial Intelligence in Special Education: Balancing Innovation and Inclusion. *International Journal of Special Education*, 41(2S), 219–233. Scopus.
- Al-Naser, Y., Sharma, S., Niure, K., Ibach, K., & Yong-Hing, C. J. (2026). Radiology Synthetic Confusion: How Generative Artificial Intelligence Amplifies Misunderstandings of Radiologists and Technologists in Patient-Facing Media. *Canadian Association of Radiologists Journal*, 77(1), 98–106. Scopus. <https://doi.org/10.1177/08465371251350085>
- Anthis, Z., & Kyriakidou-Zacharoudiou, A. (2026). Face value: How avatar identity shapes epistemic trust in AI-mediated learning. *Computers and Education: Artificial Intelligence*, 10. Scopus. <https://doi.org/10.1016/j.caeai.2026.100610>
- Arya, S., & Lin, W. (2026). Non-linear model-based sequential decision-making in agriculture. *Environmental Research Communications*, 8(3). Scopus. <https://doi.org/10.1088/2515-7620/ae489e>
- Babker, A. M. (2026). Clinical Safety and Legal Accountability of Artificial Intelligence in Hematology Diagnostic Laboratories: Risks, Failure Modes, Medicolegal Challenges, and Budgeting for

-
- Implementation. *Journal of Applied Hematology*, 17(2), 85–92. Scopus. <https://doi.org/10.4103/joah.joah.33.26>
- Balarabe, K. (2026). Algorithmic authoritarianism: Artificial intelligence's threat to privacy and freedom in the global south. *Information and Communications Technology Law*, 35(2), 201–233. Scopus. <https://doi.org/10.1080/13600834.2025.2541124>
- Balsano, C., Cabitza, F., Cicco, S., Gori, M., Malerba, D., Montagna, M., Tarquini, R., & Vacca, A. (2026). Artificial intelligence in medicine: A position paper by the Italian Society of Internal Medicine. *Internal and Emergency Medicine*, 21(1), 1–14. Scopus. <https://doi.org/10.1007/s11739-025-04146-4>
- Bergström, L. H. (2026). Ethical Issues in Artificial Intelligence and Automation. *Cineforum*, 66(2), 340–344. Scopus.
- Bougriche, Z., Osuchowski, J., Dąbroś, M., & Tomaszewski, M. (2026). Human and AI Perspectives: Comparing Emotion Recognition from Facial Expressions and Hand Gestures. *Cognitive Computation*, 18(1). Scopus. <https://doi.org/10.1007/s12559-026-10618-2>
- Ceallaigh, T. J. Ó., & Murphy, S. (2026). Teaching the teachers: A systematic review of genAI-specific technological pedagogical knowledge (TPK) in teacher education. *Computers and Education Open*, 10. Scopus. <https://doi.org/10.1016/j.caeo.2026.100367>
- Cecchi, R., Calabrò, F., Camatti, J., Santunione, A. L., Sperti, M., Zizzi, E. A., & Deriu, M. A. (2026). Artificial intelligence in healthcare: Proposal for a new medico-legal methodology in medical liability. *Legal Medicine*, 80. Scopus. <https://doi.org/10.1016/j.legalmed.2025.102764>
- Certo, L., Alves, M., & Magalhães, T. (2026). The Application of Artificial Intelligence in Public Health Surveillance in Portugal: An Exploratory Study of Expert Perspectives. *Portuguese Journal of Public Health*, 1–20. Scopus. <https://doi.org/10.1159/000550578>
- Chen, L., Zhou, J., Yao, N., Xi, H., Zhang, J., Cen, Z., Li, J., Song, Z., Li, P., Gu, Z., Baklanov, A., Christensen, J. H., Coll, I., Sahyoun, M., & Yu, S. (2026). Integrating the smart-city-brain framework for carbon neutrality, resilient and sustainable cities. *Journal of Urban Management*. Scopus. <https://doi.org/10.1016/j.jum.2026.04.001>
- Choi, A. E., Chung, M. E., & Sun, L. Y. (2026). Unmasking bias in artificial intelligence: Sex and racial representation in critical care medicine through text-to-image generators. *BJA Open*, 19. Scopus. <https://doi.org/10.1016/j.bjao.2026.100562>
- Chu, W. W., & Lam, J. L. C. (2026). Transforming healthcare management: The impact of artificial intelligence on leadership and operations. *Management in Healthcare*, 10(3), 253–265. Scopus. <https://doi.org/10.69554/AQMU2524>
- Coleman, O. F., & Waterfield, D. A. (2026). Ethical AI Use in IEP Development: A Guiding Framework. *Journal of Special Education Technology*. Scopus. <https://doi.org/10.1177/01626434261419099>
- De Gagne, J. C., Carrejo, D. J., & Hwang, H. (2026). Generative AI Ethics Toolkit for Clinical Prioritization in Nursing Education. *Nursing Ethics*, 33(4), 1169–1180. Scopus. <https://doi.org/10.1177/09697330261428627>
- del Rey Puech, P., Payne, R., Saund, J., & McKee, M. (2026). Mind the (widening) gap: Why public health must engage with AI now. *Public Health*, 250. Scopus. <https://doi.org/10.1016/j.puhe.2025.106047>
- Difonzo, F., Holst, M., Kimiaei, M., Kungurtsev, V., & Qiu, S. (2026). Towards real time control of water engineering with nonlinear hyperbolic partial differential equations. *Journal of Computational Science*, 100. Scopus. <https://doi.org/10.1016/j.jocs.2026.102945>
- Dixit, M. K., & Shanbhag, S. S. (2026). The evolution of inverse optimization: A cross-disciplinary review of algorithmic approaches and application framework in building life cycle assessment. *Energy and Buildings*, 368. Scopus. <https://doi.org/10.1016/j.enbuild.2026.117827>
-

- Emma, O. (2026). Extended reality in the digital age: A literature review on technology-driven learning environments. *Computers and Education: X Reality*, 8. Scopus. <https://doi.org/10.1016/j.cexr.2025.100127>
- Gimenez Lynch, C., & Thaller, S. (2026). Global Disparities and Artificial Intelligence in Plastic Surgery: A Narrative Review of Current Applications and Ethical Implications. *Journal of Craniofacial Surgery*, 37(3/4), 787–790. Scopus. <https://doi.org/10.1097/SCS.00000000000012054>
- Gules-Guctas, E. (2026). How Do Algorithmic Decision-Making Systems Used in Public Benefits Determinations Fail? Insights From Legal Challenges. *Public Administration Review*, 86(2), 371–382. Scopus. <https://doi.org/10.1111/puar.70043>
- Hoffarth, M. (2026). Reimagining critical thinking in the age of AI: Implications for nursing education and practice. *Nurse Education Today*, 162. Scopus. <https://doi.org/10.1016/j.nedt.2026.107076>
- Kladko, S., Lavazza, A., Yu, X., & Farina, M. (2026). From guidelines to practice: Navigating ethical and cultural challenges in operationalizing AI for project management. *International Journal of Innovation Studies*, 10(3). Scopus. <https://doi.org/10.1016/j.ijis.2026.100181>
- Kong, K., Liu, Z., Gao, Z., Dong, H., & Isleem, H. F. (2026). Hierarchical deep reinforcement learning for real-time tactical decision-making in team sports. *Data Technologies and Applications*, 60(2), 221–256. Scopus. <https://doi.org/10.1108/DTA-06-2025-0512>
- Kyambade, M., & Namatovu, A. (2026). Exploring HR Professionals' Experiences with Generative AI in Recruitment: A Phenomenological Study in Public Hospitals in Uganda. *SAGE Open*, 16(2). Scopus. <https://doi.org/10.1177/21582440261457060>
- Liu, R., & Glave, C. Z. (2026). Justice, Equity, Diversity, and Inclusion in Artificial Intelligence and Childhood Education: A Critical Systematic Review. *Journal of Research in Childhood Education*, 40(1), 5–23. Scopus. <https://doi.org/10.1080/02568543.2025.2578526>
- Madleňák, R., Madleňáková, L., Cvacho, V., & Gachulínek, D. (2026). Ethical Challenges of Artificial Intelligence in Higher Education: A Four-Pillar Student-Activity Framework for Institutional Governance. *Education Sciences*, 16(4). Scopus. <https://doi.org/10.3390/educsci16040555>
- Mariyono, D., Yunus, M., Junaidi, J., Syam, N., & Mazhabi, Z. (2026). Educational leadership in the digital era: Bridging global disparities with inclusive management strategies. *Educational Management Administration and Leadership*. Scopus. <https://doi.org/10.1177/17411432261419475>
- Massouti, A., Baroudi, S., & Shaya, N. (2026). Exploring AI integration in United Arab Emirates' teacher education programmes: The voice of educators leading the change. *Cambridge Journal of Education*, 56(3), 383–401. Scopus. <https://doi.org/10.1080/0305764X.2026.2651759>
- McConvey, K., Ghai, M., Lee, R., & Guha, S. (2026). Risk, Retention, and the Algorithmic Institution: Artificial Intelligence as a Policy Response to Higher Education in Crisis. *Canadian Public Policy*, 52(S1), 156–169. Scopus. <https://doi.org/10.3138/cpp.2025-030>
- Mo, F., Yang, Y., & Olenchak, F. R. R. (2026). Navigating the Human-AI Partnership in the Undergraduate Creative Process: A Qualitative Inquiry. *Technology, Knowledge and Learning*. Scopus. <https://doi.org/10.1007/s10758-026-09966-7>
- Molu, B. (2026). Nursing students' perceptions of gender and race in AI healthcare imagery. *Nursing Ethics*, 33(3), 813–830. Scopus. <https://doi.org/10.1177/09697330251403134>
- Muro, A. V. (2026). Theoretical pedagogical model for the integration of AI and ALFIN in the training of normal school teachers. *Revista Interamericana de Bibliotecología*, 49(2). Scopus. <https://doi.org/10.17533/udea.rib.v49n2e362266>
- Osebor, I. M., & Amaechi, E. (2026). THE ETHICAL ROLE OF NATIVE-CENTRIC ALGORITHM DECOLONIZING ALGORITHMIC BIAS IN ARTIFICIAL INTELLIGENCE. *Ramon Llull Journal of Applied Ethics*, (17), 235–253. Scopus. <https://doi.org/10.60940/rljae17Id980000016860>

- Pigac, T. (2026). Transparency in large language model (LLM)-powered digital human twins: The AI ethics perspective. *AI and Society*, 41(3), 2109–2118. Scopus. <https://doi.org/10.1007/s00146-025-02617-y>
- Qiu, M., Wu, Y., Zeng, L., & Zhang, F. (2026). Ethical considerations of AI integration in palliative care: A qualitative study. *BMC Palliative Care*, 25(1). Scopus. <https://doi.org/10.1186/s12904-026-02059-3>
- Radu, T.-A., Radu, C.-C., Pertz, C.-M., & Pașcan, E. M. (2026). ETHICAL ASPECTS IN DENTISTRY: FROM CONVENTIONAL PRACTICE TO THE DIGITAL ERA. *Romanian Journal of Legal Medicine*, 34(2), 113–119. Scopus. <https://doi.org/10.4323/rjlm.2026.113>
- Rahmatullah, T., & Nozlan, N. N. (2026). Exploring AI Integration in Financial Institutions: Challenges, Compliance, and Quality Assurance's Role in Security. *Journal of Central Banking Law and Institutions*, 5(1), 175–204. Scopus. <https://doi.org/10.21098/jcli.v5i1.445>
- Reddy, Y. J., Kurt, H., Tiwari, A., Kaur, U., Gautam Lade, S., & Waghmare, V. (2026). Optimization of IOT-Based Energy Management Systems Using Machine Learning Algorithms in Smart Buildings. *International Journal of Drug Delivery Technology*, 16(12), 721–729. Scopus. <https://doi.org/10.25258/ijddt.16.12s.86>
- Sanusi, S., Rusnawati, R., Pelas, M. U., Halik, H., & Maksum, H. (2026). Sociology of citizenship and AI ethics: Reshaping higher education amid technological disruption, a conceptual study. *Qualitative Research Journal*, 1–20. Scopus. <https://doi.org/10.1108/QRJ-09-2024-0219>
- Saroughi, M., & Nozari, H. (2026). Quantify uncertainty of meta-heuristic algorithms in optimal reservoir operation. *Journal of Hydro-Environment Research*, 64. Scopus. <https://doi.org/10.1016/j.jher.2026.100696>
- Sassenberg, A.-M., Acharya, N., Kar, P., & Eshaghi, M. S. (2026). Beyond Accuracy: The Cognitive Economy of Trust and Absorption in the Adoption of AI-Generated Forecasts. *Forecasting*, 8(1). Scopus. <https://doi.org/10.3390/forecast8010008>
- Schilhabel, S. A., Sankaranarayanan, B., Subedi, M., & Muraski, J. (2026). Exploring the impact of artificial intelligence on psychological well-being and inclusivity in higher education. *Journal of Information, Communication and Ethics in Society*, 1–17. Scopus. <https://doi.org/10.1108/JICES-09-2025-0238>
- Shal, T., Ghamrawi, N., Ghamrawi, N. A. R., Abu-Tineh, A. M., & Alshaboul, Y. M. (2026). Mobile Professional Development for STEM Teachers Using AI. *Computers in the Schools*. Scopus. <https://doi.org/10.1080/07380569.2026.2622712>
- Solano, C., Tarazona, N., Pietropaolo, A., Somani, B., & Traxer, O. (2026). Generative artificial intelligence in the daily urology unit: A practical, patient-centered framework. *Journal of Clinical Urology*, 19(4), 404–410. Scopus. <https://doi.org/10.1177/20514158261445149>
- Sopiyan, M., Judijanto, L., Asad, M., Zuhri, S., & Syufa'at, S. (2026). Artificial Intelligence in Islamic Family Law: Addressing Ethical and Regulatory Gaps through a Maqāṣid al-Sharī'ah Approach. *Mawaddah: Jurnal Hukum Keluarga Islam*, 4(1), 1–18. Scopus. <https://doi.org/10.52496/mjhki.v4i1.31>
- Sullivan, A., Geis, G., & Cummings, C. (2026). Ethics Education in Neonatology: Integrating Theory, Multimodal Methods, and AI Innovation. *NeoReviews*, 27(2), e73–e83. Scopus. <https://doi.org/10.1542/neo.27-2-099>
- Tella, A., & Olatoye, A. A. (2026). Artificial intelligence-driven gamification in smart libraries for redefining user engagement in the library of the future. *Library Hi Tech News*, 43(5), 4–8. Scopus. <https://doi.org/10.1108/LHTN-04-2026-0095>
- Tian, L., Wu, W., Tan, Y., & Feng, S. (2026). Self-reported AI literacy and gendered self-efficacy: Evidence from Chinese adolescents. *Telecommunications Policy*, 50(4). Scopus. <https://doi.org/10.1016/j.telpol.2026.103152>

- Uchoa, A. P., Oliveira, C. E. T., Motta, C. L. R., & Schneider, D. (2026). Natural-Language Mediation Versus Numerical Aggregation in Multi-Stakeholder AI Governance: Capability Boundaries and Architectural Requirements. *Computers*, 15(1). Scopus. <https://doi.org/10.3390/computers15010024>
- Wright, G. W., Diaz, B., & von Gillern, S. (2026). Science Educators' Attitudes and Perspectives on Artificial Intelligence (SEAP-AI) Scale. *Journal of Science Education and Technology*. Scopus. <https://doi.org/10.1007/s10956-026-10302-y>
- Yadav, A. B. (2026). Strengthening consumer consent in e-commerce: Legal and policy reforms to address dark patterns and AI-driven challenges. *Journal of Data Protection and Privacy*, 8(3), 254–271. Scopus. <https://doi.org/10.69554/MYRT2092>
- Zhai, X. (2026). What a party knows: Rethinking knowledge assessment in automated contracting. *Computer Law and Security Review*, 61. Scopus. <https://doi.org/10.1016/j.clsr.2026.106336>
-

Copyright Holder :

© Yayu Laode et.al (2026).

First Publication Right :

© Al-Hijr: Journal of Adulearn World

This article is under:

