

The Ethics of AI in Financial Planning: Bias, Transparency, and the Role of Human Judgment

Dolapo Achimugu¹, Chinaza Ukatu¹, Arinze E. Anaeye²

¹College of William & Mary, Raymond A Mason School of Business, Williamsburg, Virginia US

²Department of Accounting, Kingsley Ozumba Mbadiwe University, Ideato, Nigeria

DOI: <https://doi.org/10.51244/IJRSI.2025.1213CS001>

Received: 10 September 2025; Accepted: 16 September 2025; Published: 15 October 2025

ABSTRACT

The fast-growing use of artificial intelligence (AI) has introduced new ethical issues in the financial services sector. Robo-advisors, loan algorithms, and automated financial instruments now make choices that impact people's lives substantially. These automated instruments may lack fairness, clarity, and human supervision. Without adequate checks, they could generate discriminatory decisions or erode trust in financial institutions. This paper sets forth a normative-ethical framework to help oversee the responsible use of AI in financial planning. The study identified a novel framework known as the EFT Model, which has four pillars: Ethical Design, Fairness, Transparency, and Human Oversight. The paper examines each principle in detail, illustrating it with practical examples like discriminatory loan approvals and unclear investment recommendations. Roles and accountability of key players such as developers, regulators, financial institutions, and customers are also clearly identified. The paper harmonizes the framework with existing regulations like the EU AI Act, the GDPR and discusses how it could help direct ethical design in practice. It also underlines the importance of conducting additional research with the intention of testing and refining the model under real-world conditions.

Keywords: AI Ethics, Fintech, Ethical Framework, Human Oversight, Transparency, Fairness

INTRODUCTION

Financial services are being rapidly transformed with the use of Artificial Intelligence (AI). Among the most significant applications to financial planning are AI-enabled products such as robo-advisors, scoring engines, and algorithmic investment products (Sahu, 2024). These hold the promise of higher efficiency, scale, and differentiated service offerings, often beyond human capacity in aggregating vast sets of financial information. As investors and institutions rely increasingly on AI support, not only in investment decisions, but in insurance planning, retirement, and in wealth management, its effects on financial outcomes keep growing (Vuković et al., 2025). The rapidity with which AI has been introduced into financial planning, however, poses ethical issues. Even as efficiency enhancement and cost reduction are normally the key driving forces underlying AI introduction, bias, lack of explainability, and lack of human supervision are emerging risks (Černevičienė & Kabašinskas, 2024). AI algorithms utilized in scoring, such as those employed in underwriting, are found with tendencies to perpetuate existing social biases (Agarwal, 2024). In some cases, customers' loan applications are rejected or offered at unfavorable terms. The lack of explainability limits the autonomy of the customer, reduces trust, and creates gaps in accountability (Nallakaruppan et al., 2024).

Gladstone and Hundtofte (2023) noted that financial planning is an area where decisions make long-duration impacts on people's financial well-being and safety. It involves a great deal of trust, ethical judgment, and contextual awareness. These traits are not always transferable to AI systems, particularly black-box AI like deep learning. The lesser role played by human advisors at key points in decisions has additionally raised an ethical issue about offloading moral responsibility onto machines. While AI can facilitate decisions, the elimination of human intervention in decisions involving high financial stakes becomes ethically problematic (Giarmoleo et al., 2024). Despite rising awareness about these challenges, most of what has been written about AI in financial

settings tends to center around the technical solutions like bias detection programs, explainable AI tools, and compliance structures (Hermosilla et al., 2025; Saarela & Podgorelec, 2024). AI ethics in financial planning remains fragmented; existing frameworks (e.g., OECD, IEEE) are general and lack sector-specific guidance. While regulations like GDPR or Basel III offer compliance structures, they do not address the moral reasoning needed for high-stakes decisions like retirement planning or risk profiling. Procedural tools (e.g., explainable AI, fairness audits) support transparency but fall short of defining what ethically ought to be done. While significant, these innovations fail to adequately confront the underlying ethics about what financial AI systems should do, who should be held accountable, and what fairness should mean in financial decisions. No unified normative ethics framework currently exists in this area that should inform AI's design, deployment, and governance in financial planning settings. Morley et al. (2021) suggest that most AI ethics frameworks fail to provide domain-specific guidance. Also, Jobin et al. (2019) reviewed 84 documents containing ethical principles or guidelines for AI and found that no single, unified, or enforceable ethical framework exists across sectors.

This paper fills this crucial gap. It makes the case for a normative framework of ethics built around three key issues: bias, transparency, and the use of human judgment. Rather than defining the issue technically or legally, this paper borrows ideas from normative ethics, particularly theories of fairness, responsibility, and moral agency, to make the case for an ethical approach to AI in financial planning. The objective goes beyond risk identification to provide an ethically oriented direction for developers, financial firms, regulators, and other players in the financial technology sector. The paper's contribution has two parts. Firstly, it presents a theoretically informed ethical framework that incorporates fairness, transparency, and human agency into AI-informed financial decisions. The framework has its roots in normative ethics and theory, but takes the form that might inform practical use. Secondly, the paper considers how this framework might apply in real-world financial situations, like automated investment recommendations or credit scoring.

Conceptual Foundations

AI in Financial Planning

Artificial Intelligence (AI) in financial planning refers to the use of computer programs that are capable of conducting tasks that would normally involve human intelligence. Such tasks are pattern recognition, decision-making, and forecasting (Najem et al., 2025). In financial services, AI programs are continuously incorporated into services with the view of boosting precision, speed, and customization (Vuković et al., 2025). There are two categories of AI applications in financial planning. The first is decision support systems. These systems assist human financial advisors by providing advanced analytics and recommendations. They do not make the final decision but offer data-driven insights. Examples include risk analytics tools and market trend predictors. The other category includes decision automation systems. These are automated systems that make decisions without direct human intervention. Robo-advisors are an obvious example. They apply algorithms to make investment portfolio recommendations and adjustments, given user choices and market trends (Jia et al., 2022; Tao et al., 2021). The other area where AI is employed includes credit scoring (Raji et al., 2024). The conventional credit scoring approaches depend on fixed indicators such as repayment history and income. AI, however, employs machine learning algorithms that deal with large, diverse, but relevantly disparate sets of information, including social media and history of transactions. The systems are capable of arriving at more dynamic, inclusive scores (Li et al., 2024). That, however, creates fairness issues, particularly when the information sources mirror past discriminatory trends.

Portfolio management represents another key use. AI-driven robo-advisors like Betterment or Wealthfront make automated decisions about asset allocation with respect to market conditions and risk profiles (Lam, 2016). Products like these make financial advice cheaper and more accessible. However, they also eliminate the human factor in difficult financial choices, potentially impacting ways that compromise client trust and emotional comfort (Ahmad et al., 2023). AI also finds use in risk management. Algorithms track volatility in the market, flag potential instances of fraud, and evaluate systemic risk. Institutional financial use includes AI models enabling stress scenarios for simulations, along with predicting economic trends. These, in turn, run typically in real time, providing current information to decision-makers.

However, when these instruments are utilized with non-transparent models or without regulation, they can generate values that users cannot question or explain (Bahoo et al., 2024). It is useful at this point to differentiate narrow AI from general AI. Today's financial programs are primarily narrow AI, which refers to a specific use program. These are effective under defined constraints but do not possess general capabilities in giving advice, in the form of general reasoning (Walton, 2018). This renders them effective, yet narrow, particularly where unpredictable, new situations arise. AI in financial planning covers an enormous range of products and capabilities. These range from aiding human advisors to complete automated processes. While possessing advantages like effectiveness, accessibility, and speed in computing, they also pose an additional set of risks, both practical and ethical. It is crucial, before determining its normative effects, to understand where AI becomes incorporated into financial structures.

Ethics in Technology Use

Ethics in relation to technology, particularly Artificial Intelligence (AI), becomes more relevant as these programs make decisions that have significant impacts on people's lives. The ability to assess the morality of these decisions requires an education in ethical theory. Ethics involves the examination of right and wrong, and how people should act. There are various branches within this area. Two important ones are normative ethics and applied ethics (Chaddha & Agrawal, 2023). Normative ethics involves attempting to create theories and principles, informing what people should do. It raises questions such as "What is the right thing to do here?" and "What's our duty to others?" It does not become involved with specific situations but rather attempts to create abstract rules and moral guidelines. For instance, it could investigate whether fairness should guide decisions or whether decisions should maximize overall utility (Dempsey et al., 2023). By contrast, applied ethics takes these theories and attempts to apply them in relation to specific, practical issues. Applied ethics is inherently concerned with how AI systems impact human beings. It focuses on how the systems are developed, the processes, logic of decision making, who makes decisions, how and to what extent. These range from controversies over data privacy, algorithmic discrimination, or AI opacity (Bleher & Braun, 2023; Kazim & Koshiyama, 2021). In this way, normative ethics lays the foundations, while applied ethics brings it into practice.

Several normative ethical theories are useful for understanding and critiquing AI systems in financial planning. The three most relevant are deontology, utilitarianism, and virtue ethics. Each offers a different way to judge the morality of AI development and use.

Deontology, as represented by scholars such as Immanuel Kant, stresses duty and rules. In deontological thought, one fundamental notion holds that certain actions are morally obligatory or prohibited, regardless of the consequences (Barrow & Khandhar, 2023). For instance, when an AI tool employed in the field of scoring arrives at correct conclusions but discriminates against some persons, it remains unethical, according to the deontologist. This is because it violates a moral duty to treat people equally and respect their rights. Deontology also emphasizes transparency and accountability. Users and regulators should be able to know how decisions are made. If an AI system cannot explain its decisions or allow individuals to contest them, it may be seen as morally unacceptable under this view (Jedličková, 2024).

Utilitarianism, on the other hand, is focused on outcomes. It holds that the best action is the one that produces the greatest good for the greatest number. When applied to AI in financial services, this theory looks at whether an algorithm improves financial access, reduces costs, or benefits more people than it harms (Anshari et al., 2022). A utilitarian might support a system that increases efficiency and reduces overall bias, even if a small number of individuals are negatively affected. However, this approach can sometimes justify unfair treatment of minorities if it benefits the majority. This tension raises ethical concerns in financial contexts, especially where long-term inequalities may be reinforced by AI models trained on biased data (Card & Smith, 2020).

Virtue ethics takes a different approach. It does not focus on rules or outcomes but on the character of the people and institutions involved. This theory asks whether the design, development, and use of AI reflect virtues like responsibility, honesty, and integrity (Hagendorff, 2022). In the financial sector, it could be whether developers are cautious when they are training models or whether financial firms are honest and transparent with users. Virtue ethics encourages an ethics-aware culture, not technical compliance alone. It also prefers the approach of

responsible innovation, where ethical thinking becomes part of innovation at the onset, not an afterthought (Griffin et al., 2024).

Alongside traditional theories, there are also AI-specific ethical frameworks. These frameworks are designed to address the unique features of AI systems, such as autonomy, opacity, and data dependency. One of the most influential is the work of Luciano Floridi and colleagues, who developed principles like non-maleficence (do no harm), beneficence, justice, and explicability (Floridi & Cowls, 2019). These principles combine elements of normative ethics to guide the development of trustworthy AI.

Another milestone is the European Union's AI Act, projecting risk-based regulation of AI applications. In April 2021, the European Commission proposed the first EU regulatory framework for AI, which was later publicized as Regulation (EU) 2024/1689 (Regulation (EU) 2024) on 12 July 2024, establishing a governance structure and setting out clear requirements for the Commission and the AI Office. It classifies AI systems in several risk categories, ranging from minimal to unacceptable. The higher-risk ones, such as those used in employment choices or credit ratings, face more stringent regulations. It aims at offering fairness, human oversight, and explainability in automated decisions (Szadeczky & Bederna, 2025). The IEEE P7000 series is another regulation. It was launched by the Institute of Electrical and Electronics Engineers (IEEE) as a set of standards committed to embedding ethically relevant considerations into AI system building. For example, the IEEE P7001 standard touches on transparency, with the provision that users should be able to understand and question decisions made by intelligent systems (Spiekermann, 2017). These guidelines often overlap with mainstream ethical theories but are set down with technological deployments in focus.

These frameworks and theories lay a strong base for analyzing the ethics of AI in financial planning. They guide not only what AI can do, but what it ought to do and under what circumstances. No one theory has all the answers, but together, they permit an enriched, better-balanced ethical evaluation.

Key Ethical Tensions in Financial AI Systems

AI technologies in financial services hold several ethical tensions caused by the intersection between efficiency, fairness, accountability, and autonomy. These tensions are not just technological difficulties but ethical challenges that impact individual rights, institutional trust, as well as societal equity. Handling these issues involves an understanding of the nature of AI systems and ways in which their use in financial settings creates tensions that are ethically problematic.

Bias and Discrimination

Hanna et al. (2024) indicate that AI bias represents a key ethical issue. AI financial systems frequently use past data to train algorithms. If such data includes past inequality or systematic discrimination, an AI system could reproduce or even magnify bias. For example, the algorithms for scoring credit could act against minority populations because of ingrained patterns in the training set. Even neutral-looking variables such as zip codes or school names become proxies for race or socio-economic background, causing discriminatory decisions about giving credit (Cristina et al., 2023). The dilemma lies between AI's promise of efficiency and fairness demanded in financial services. Companies may try to maximize prediction and outcomes, but this optimizes at the expense of treating people fairly. Biases in algorithms could pass performance tests but consistently discriminate against populations. The dilemma points to the limitation of exclusively data-based systems and the importance of ethical regulation.

Transparency vs. Complexity

Most AI applied in finance, like deep learning models, act as "black boxes." The inner mechanisms are not clear, even to experts. This unpredictability causes an important ethical tension (Svetlova, 2022). Financial judgments involve high stakes, like granting loans, distributing investments, and setting insurance rates. However, users and affected persons frequently cannot see how to review the rationale behind decisions. This tension places model precision and sophistication against the ethical need for explainability. Regulators and ethicists maintain

that people should be entitled to comprehend decisions that impact their financial well-being. The EU's General Data Protection Regulation (GDPR) has a "right to explanation," but it is difficult in practice when AI systems are "black boxes" (Wachter et al., 2018). This also causes an issue with "information asymmetry." AI builders and financial firms are normally better versed technically than their customers, causing uneven power dynamics. Without clarity, users cannot provide informed consent or question decisions, eroding autonomy and accountability.

Automation vs. Human Judgment

AI allows high levels of financial decision-making automation. Robo-advisors and automated loan processors keep prices low and ensure 24/7 operations. There is an efficiency cost, however, in that human judgment is not present (Maier et al., 2022). AI has no compassion, moral sense, or capacity to consider individual situations. The outcome can be callous or rigid decisions. There is an ethical trade-off between speed/consistency and contextual sensitivity/moral judgment (Farisco et al., 2020). For example, a human loan officer might waive an unfavorable credit report because of strong individual circumstances. An AI program might lack that flexibility. The tension is particularly relevant where financial distress intersects with sensitive populations. Also, AI dependence has the potential to induce "automation bias," where individuals follow algorithmic counsel even when they suspect errors. It erodes the role of professional judgment in financial matters and obliterates important checks and balances. The use of human-in-the-loop mechanisms has been proposed as a solution, but it creates its own set of ethics around responsibility and liability (Salloch & Eriksen, 2024).

Data Privacy and Surveillance

Financial AI platforms deal with enormous quantities of personal information, such as income, consumption patterns, credit history, and even social media behavior. While this information increases predictive ability, it also creates questions about consent and privacy (Aldboush & Ferdous, 2023). The users do not know what happens to their information, who accesses it, or how long it remains in storage. The ethical dilemma inherent in this situation involves striking an equilibrium among data effectiveness in driving innovation, risk evaluation, and individual rights over privacy, as well as data safeguard. Financial institutions can contend that data-led personalization has advantages. However, without clear guidelines, information gathering becomes overbearing and compulsory. AI-powered monitoring could result in profiling and manipulation, causing mistrust in financial products. The dilemma gets compounded by governance gaps. While some places, such as the EU, possess effective data safeguard regulations, others lack effective frameworks. Incompatibilities in governance over information leave cross-border financial multinationals with an ethical dilemma, as they operate across borders (Thein et al., 2024).

Accountability and Moral Responsibility

Machado et al. (2024) noted that one of the ethical conflicts in AI-based finance is accountability. When an algorithm causes an adverse or discriminatory outcome, it is unclear what entity, if any, should be held accountable—the programmer, the bank, the source of information, or the program. This diffusion compromises redress with an ethical basis. In conventional financial frameworks, chains of accountability are clear. With AI, responsibility becomes fragmented. This leads to what is sometimes called the "moral crumple zone," where human actors absorb blame for decisions made by opaque systems (Kaas, 2024). This tension threatens both legal clarity and moral justice. There is growing consensus that ethical AI systems in finance must include accountability frameworks—clear documentation, audit trails, and transparent governance structures. Without these, financial AI remains a system where errors are difficult to trace and ethical violations are easy to deflect (Cheong, 2024; Raji et al., 2020). These tensions are visually summarised in Figure 1, which illustrates the ethical trade-offs that frequently emerge in the design and deployment of AI systems in financial planning.

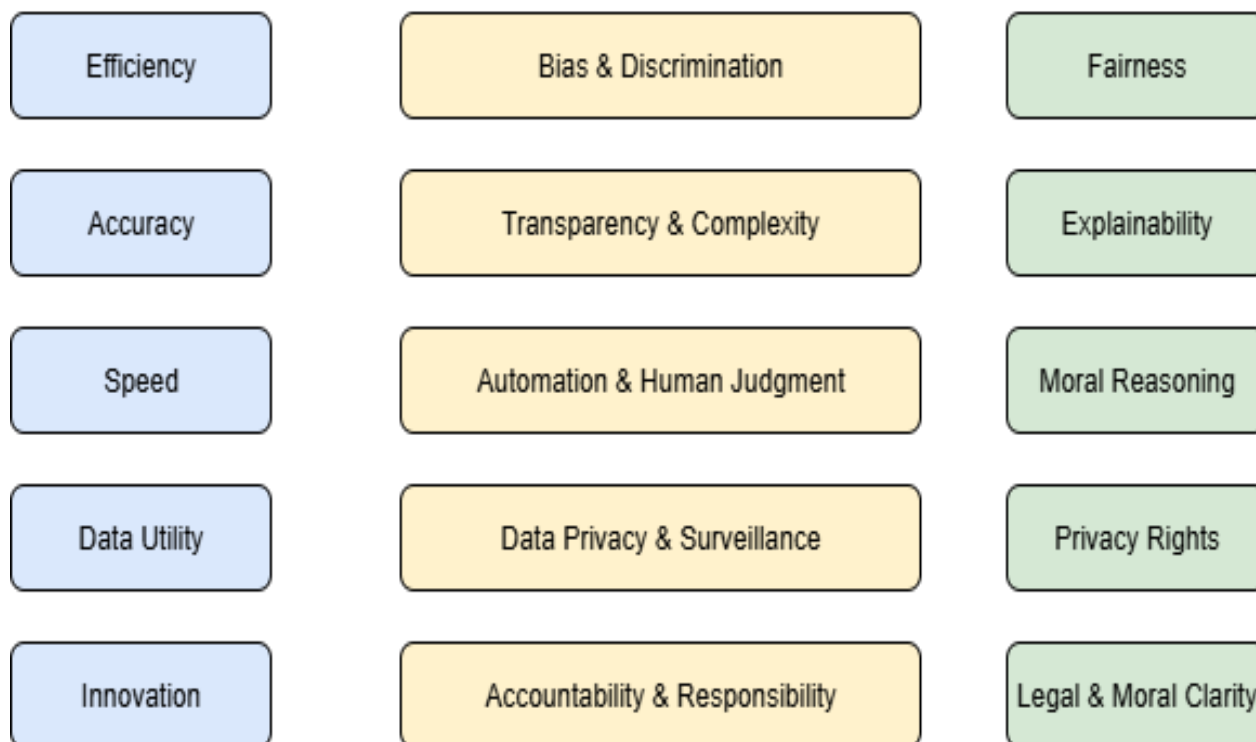


Figure 1: Key Ethical Tensions in Financial AI Systems

Core Ethical Issues in AI-Driven Financial Planning

Algorithmic Bias

Algorithmic bias refers to systematic and repeatable errors in AI outputs that unfairly disadvantage certain groups (Ukanwa, 2024). In financial planning, such biases can influence important decisions like loan approvals, credit scoring, investment risk profiling, and insurance pricing. While these systems are often marketed as objective and data-driven, they can perpetuate or even amplify societal inequalities embedded in their training data or decision rules (Cristina et al., 2023; Fuster et al., 2021). Bias can enter AI systems through several pathways. One common source is the training data. If historical financial data reflects discriminatory lending practices or underrepresentation of certain populations, an AI model trained on that data will likely learn and reproduce those same patterns (de Castro Vieira et al., 2025; Nwafor et al., 2024). Feature selection also plays a critical role. Even if sensitive variables like race or gender are removed, proxies such as ZIP codes, education level, or employment history can indirectly encode discriminatory patterns (Wang et al., 2024). Additionally, biased assumptions built into model architecture or optimization criteria (such as maximizing accuracy over fairness) can further entrench inequality.

There are growing examples of algorithmic bias in real-world financial services. One well-known case involved a major credit card company whose algorithm gave significantly lower credit limits to women than men, even when both had similar financial profiles. In another case, a digital lending platform disproportionately rejected minority applicants, despite their creditworthiness. Such disparities highlight how automated decision systems, if left unchecked, can perpetuate racial and gender discrimination under the guise of efficiency and neutrality (Cristina et al., 2023). From an ethical perspective, algorithmic bias raises concerns about justice, fairness, and non-discrimination. Deontological ethics emphasizes respect for individuals' rights and equal treatment under rules. Under this view, an AI system that treats similar individuals differently based on irrelevant or prejudicial factors violates the moral duty of fairness (Ebrahimi et al., 2024). Even if biased outcomes result in economic efficiency (e.g., by targeting the most "profitable" borrowers), a utilitarian approach must weigh these benefits against the broader social harms and loss of trust in financial institutions. The injustice suffered by individuals unfairly denied access to financial opportunities cannot be justified by aggregate economic gain.

Virtue ethics further emphasizes the moral character of decision-makers and institutions. Developers and financial professionals who deploy AI systems bear a responsibility to build and use them with care, integrity, and a commitment to social justice. Algorithmic fairness should not be an afterthought or optional feature; it is a core requirement of responsible innovation. These encompass the use of fairness-aware machine learning, auditing bias in systems, and engaging multi-stakeholders in the design of systems (Hagendorff, 2022). A normative approach to this question holds that AI in financial planning should be transparent, non-discriminatory, and fair. Financial institutions operate with immense power over people's economic lives. When they use AI tools, they take upon themselves the ethical responsibility of ensuring those tools do not deepen social inequalities. Technical solutions such as de-biasing algorithms and explainable AI are important but insufficient. Ethical supervision, human judgment, and accountability mechanisms should go along with technological protection (Kowald et al., 2024).

Transparency and Explainability

Financial planning AI solutions frequently use sophisticated algorithms. These are often deep learning algorithms that cannot be interpreted to explain how exactly a specific choice was made. These algorithms are sometimes referred to as “black-box” algorithms because they cannot provide transparent explanations to both financial experts and users (Černevičienė & Kabašinskas, 2024). In financial services, a lack of transparency presents severe ethical challenges. The financial decisions influence people's access to credit, investment, creditworthiness, as well as future financial stability. If an individual has been given a recommendation or choice where the rationale cannot be explained, this compromises the issue of informed consent as well as erodes trust.

The ethical appeal for explainability in financial AI systems corresponds with higher values of autonomy and accountability. From an autonomy standpoint, people should understand the rationale behind decisions made about their financial lives. By knowing, they are in a better position to decide. Without explainability, autonomy suffers. Moreover, a lack of explainability complicates the ability to catch errors or bias in the system, with implications for accountability and fairness (Wachter et al., 2017). An example in practice includes the early adoption of robo-advisors. Some customers complained about receiving portfolio recommendations that were inconsistent with risk tolerance, with the system offering no adequate explanation (Boreiko & Massarotti, 2020). In this scenario, the black box nature of the algorithm causes challenges in questioning the recommendation, not knowing if the recommendation was correct, and worst, suffering financial loss. In ethics, this is considered going beyond a technological issue. It becomes a morality issue. According to deontological ethics, people should be respected and not considered passive recipients of decisions. Explainable AI allows fairness, respects dignity, and supports the creation of accountable AI systems (D'Alessandro, 2024).

Human Judgment and Accountability

Growing dependence on AI in financial planning has also resulted in an increasing trend toward complete automation. Algorithmic decision systems and robo-advisors are now able to administer portfolios, evaluate risk, and provide investment product recommendations without human intervention (Boreiko & Massarotti, 2020). While speed and efficiency are offered, questions about errors, bias, and loss of public trust arise when human advisors' roles are eliminated. It is an important question whether the elimination of human supervision increases financial systems' vulnerability to errors, bias, or loss of public trust. Elimination of human judgment has serious risks. When errors are made, automated programs multiply small errors on a large scale. While humans are sensitive to emotions, AI lacks sensitivity as well as the ability to interpret the distinctive circumstances of the clients. It can result in technically correct but ethically inappropriate recommendations. The studies proved users feel uncomfortable dealing with fully automated financial instruments, mainly when decisions lack clarity, empathy, or understanding (Klingbeil et al., 2024; Zhu et al., 2024).

Ethical reasoning favors the continuation of human judgment. Empathy, understanding of context, and moral accountability are human capabilities that cannot be fully emulated by machines. The question of moral agency suggests that an individual has to be held responsible for decisions, particularly when causes of harm arise. It is why proposals like human-in-the-loop and significant human control have been made (Santoni de Sio & Van den Hoven, 2018). These frameworks keep humans within the decision-making continuum, with control over

what happens and the right to override when called upon. Human judgment is needed not only to correct mistakes but to maintain values like care, trust, and accountability in financial services.

Proposed Normative-Ethical Framework

Ethical Design Principles: The EFT Model

In response to the ethics issues in AI-based financial planning, this paper proposes a novel framework known as the EFT Model. This is a normative-ethical framework constructed upon four pillars: Ethical Intent, Fairness, Transparency, and Human Oversight. Each one aims at an essential area of ethics concern and has been crafted to ensure the prudent advancement as well as use of financial AI systems.

Ethical Intent entails the incorporation of values into the design process right from the beginning. AI systems should be constructed with the intention to advance social good, the welfare of the client, and professional ethics. Developers, as well as financial institutions, should give active thought to the potential damage or abuse of their systems. Ethical intent encompasses the values of beneficence and non-maleficence, commonly referenced in bioethics, and is equally applicable in financial situations (Floridi et al., 2018; Jobin et al., 2019).

Fairness demands that AI systems be designed to prevent bias and discriminatory outcomes. This includes regular audits of training data, validation of algorithms across demographic groups, and corrective actions where disparities exist. Bias can enter through data, model selection, or even developer assumptions. Therefore, fairness must be both a design goal and a regulatory requirement (Buolamwini & Gebru, 2018; Mehrabi et al., 2019).

Transparency is essential to build trust and enable accountability. Financial AI systems should disclose their methods, criteria, and logic used in decision-making. Explainable AI (XAI) techniques are crucial in helping clients understand why specific recommendations or rejections occur. Without transparency, it becomes unfeasible to obtain informed consent, and users could disengage or resist the system (Madaan, 2025).

Human Oversight ensures that critical decisions do not occur in isolation from human judgment. Human-in-the-loop systems retain a layer of interpretive control, especially for high-stakes decisions. Assigning responsibility is part of this principle, helping ensure that when errors occur, accountability is traceable (Santoni de Sio & Van den Hoven, 2018).

The EFT Model constitutes an operational manual for aligning AI instruments with financial services' ethical values. The flow and organization of the suggested Ethical Framework (EFT) are diagrammatically represented in Figure 2.

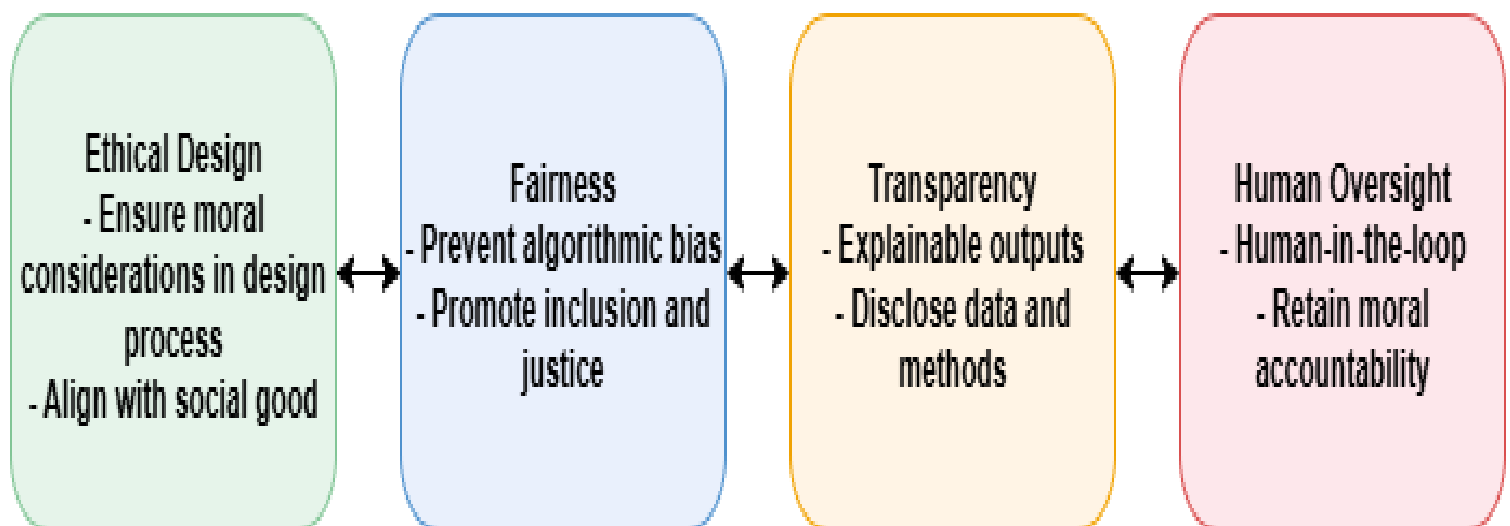


Figure 2: EFT Model for AI-Driven Financial Decision-Making

Stakeholder Roles and Obligations

In an ethically sound financial AI landscape, every major player has an active role in designing secure, just, and trustworthy systems. While their roles are distinct, they are interdependent and need to be congruent in order for the system to operate optimally.

Developers are the first point of ethical duty. They create, write, and implement algorithms, with those decisions affecting all users. It falls upon developers to make it fair, transparent, and accurate. That includes employing varied sets of data, verifying bias, and applying explainable AI models (Heidari et al., 2019). In addition to technological skills, developers should be educated in ethical thinking and innovation responsibility (Griffin et al., 2024).

Ridzuan et al. (2024) noted that financial institutions are also promoters of AI tools. They are responsible for integrating AI ethically into financial services. This means choosing vendors who comply with ethical standards, conducting regular audits of AI performance, and ensuring that clients understand how recommendations are made. Institutions must provide clear documentation and avenues for appeal when clients dispute results. They also bear ultimate accountability for harm caused by the systems they use.

Regulators act as external guardians. Their role is to create policies that guide AI use, define fairness and transparency standards, and enforce compliance. In fast-evolving fields like AI, regulators must also stay updated and adapt rules to emerging risks. Regulatory sandboxes, for instance, allow testing AI systems under supervision before full deployment (Yordanova & Bertels, 2023).

Clients also have responsibilities. As end-users, their decisions often rely on AI outputs. Therefore, digital financial literacy is essential. Clients must understand basic concepts like risk profiling, data sharing, and AI limitations. Institutions must support this by offering user-friendly tools and educational resources (Amnas et al., 2024).

Shared ethical responsibility ensures that no single group bears the burden alone. A multi-actor approach strengthens trust and accountability across the financial AI ecosystem.

Ethical Decision-Making Flow

Financial decisions can have significant consequences, especially when involving AI. A systematic process for ethical decisions reduces damage, contains risk, and preserves human control. The three-stage flowchart consists of Trigger Points, Ethical Checkpoints, and Escalation Paths.

The first level is the trigger point, where an AI financial system has encountered a high-stakes or ethically sensitive situation. These could be investment choices beyond some threshold, loan grants, or retirement account disbursements. These situations flag the system automatically for an extra level of check.

The second phase includes ethical checkpoints. These are integrated criteria that assess the recommendation's fairness level, transparency, and accuracy. For example, to what extent has the algorithm explained its recommendation adequately? Are there signs of bias, conflicting information? These checkpoints are like internal auditing.

In the event that a checkpoint fails, the process escalates, calling upon human review. It then becomes the responsibility of a financial advisor or compliance officer to intervene, reviewing the system manually. Human reviewers are trained in applying contextual judgment, understanding client worries, and making ethically sensitive choices. This phase maintains significant human control with the prevention of loss of moral agency in automation (Santoni de Sio & Van den Hoven, 2018).

This flow allows ethical vigilance without sacrificing efficiency. It can be presented in a simple decision tree diagram. The visual shows where automation operates, when ethical checks apply, and when human input is

required. Institutions can adapt this structure to fit their service types and risk levels. The flowchart in Figure 3 illustrates the ethical decision-making process embedded within financial AI systems, highlighting key checkpoints, escalation triggers, and pathways for ensuring accountability and transparency in automated high-stakes decisions.

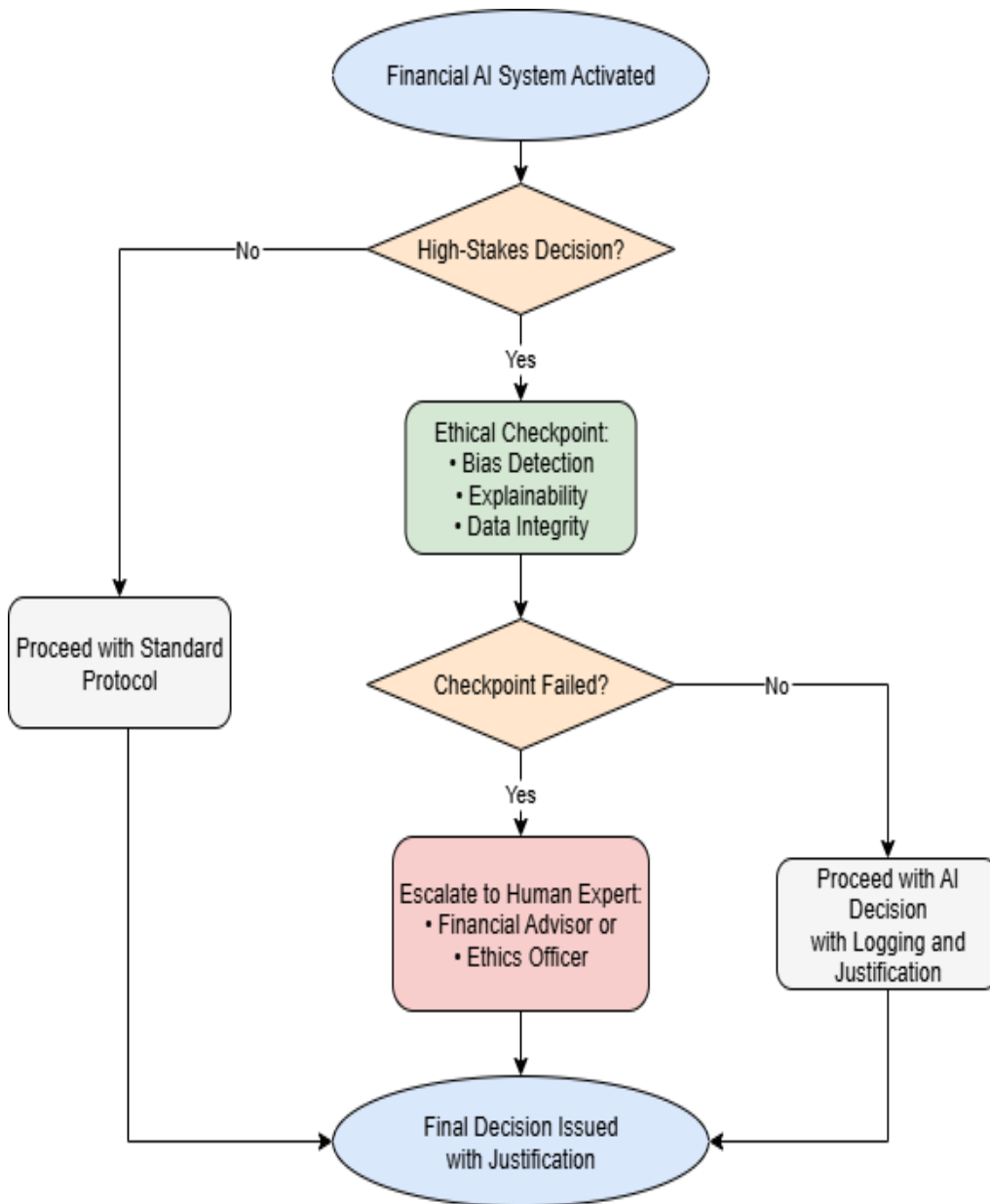


Figure 3: Ethical Decision-Making Flow for Financial AI Systems

Application of Framework to Real-World Scenarios

The EFT model can be applied to the following real-world scenarios to address the issues created by the use of AI in financial services.

Robo-advisor investment portfolio with opaque methodology

Many retail investors use robo-advisors to manage their investments. However, some of these systems do not explain how portfolio recommendations are generated. This lack of transparency can cause confusion and distrust, especially during market downturns. Applying the EFT framework, the Transparency pillar would require explainable AI (XAI) techniques, where the user is shown a clear rationale behind each investment choice. Ethical responsibility dictates that the developer include audit logs and visual summaries. The Fairness pillar

ensures that all users receive equitable risk profiles regardless of financial literacy. This level of openness can reduce panic-driven withdrawals and build user confidence.

AI loan algorithm biased against minority applicants

In 2019, Apple Card faced allegations that women were given lower credit limits than men with similar financial profiles (BBC, 2019). Similar concerns have been raised in U.S. mortgage lending data (Liu & Liang, 2025). Using the EFT framework, the Fairness component mandates proactive algorithmic audits and bias detection protocols. Transparency would require lenders to disclose how creditworthiness is calculated, especially when traditional credit scores are supplemented with alternative data. Human oversight becomes critical in high-impact decisions, such as loan approvals. It would allow flagged decisions to be reviewed manually. This helps restore public trust, as well as comply with regulatory requirements.

Automated retirement recommendations without individual context

Some financial planning websites supply retirement horizons without considering factors that may be specific to certain users, like continued illness, dependents, or late-career turbulence (Gorry & Leganza, 2024). These present ethical concerns with individualization and human dignity. Using the EFT paradigm, Ethical responsibility may involve requiring users to provide contextual variables. Fairness means requiring the system to account for life-stage variation. This keeps it relevant and prevents damage caused by one-size financial advice.

Implications

The Ethical Framework (EFT) presents practical, regulatory, and research avenues for increasing prudent AI application in financial services.

For Practice

Fintech developers are primarily responsible for implementing the EFT framework at the design level. Ethical values should be infused in system architecture. That involves utilizing varied training data, fairness testing, and explainable model use. The developers should ensure that decision rules are explainable and outcomes are interpretable. UX interfaces should also exhibit transparency with explanations and opt-out options provided for users. These measures avoid eroding trust and facilitate responsible innovation.

Financial planners should not depend exclusively on automated systems. Rather, active supervision should be preserved, particularly in high-risk decisions. The human-in-the-loop concept guarantees that clients should be able to question automated recommendations. The use of the EFT framework should help planners decide when technology should be employed and when individual judgment should be exercised. It keeps the interests of the client secure while ensuring the duty of care. System output should also be frequently reviewed to identify potential causes of harm.

For Policy and Regulation

The EFT framework aligns closely with Article 22 of the GDPR, providing individuals with the right not to be made subject to decisions under fully automated processes. Even the EU AI Act preserves financial decision-making tools as having the potential for high risk and specifies the need for transparency, fairness, and human oversight. In America, the FTC and Consumer Financial Protection Bureau (CFPB) guidelines provide similar warnings about the black box mechanisms of AI leading to discrimination. The EFT model offers usable tools to meet such legal mandates. It could also complete ISO/IEC 42001 AI management systems standards with specific ethical directions, along with procedures for safeguards.

For Research

There is a need for empirical research to confirm the EFT framework in multiple financial settings. Future research should experiment with how the framework shapes user trust, decision making, and bias minimization.

It is possible to combine the use of EFT metrics with AI auditing instruments like IBM's AI Fairness 360 or Google's What-If Tool. By measuring ethical indicators, researchers can create guidelines for responsible Fintech. It will also help contribute to the development of industry-wide certification programs.

CONCLUSION

Rising AI in financial services requires an unambiguous and consistent ethics strategy. As automation spreads, the dangers of bias, lack of transparency, and loss of human agency also escalate. There is a need for a normative ethical framework, ensuring that AI systems are compatible with fundamental human values and regulations. The proposed Ethical Framework (EFT), grounded in fairness, transparency, and human supervision, directly confronts the issues. It encourages algorithmic review, informs system logic, and enforces human accountability in decisions. These values are not only moral obligations but also practical instruments for constructing user trust and institutional legitimacy. Fair and transparent AI underpins regulation compliance while enabling clients to make independent decisions. Human agency remains fundamental, where empathy and moral judgment are important, especially in high-risk situations. This novel framework fills an important gap in current practice and suggests a direction towards accountable AI governance in financial sectors.

REFERENCES

1. Agarwal, V. (2024). Fair or Flawed? Assessing AI's Impact on Credit Decisions. *International Journal of Computer Trends and Technology*, 72(Dec 12, 2024), 128–132. <https://doi.org/10.14445/22312803/IJCTT-V72I12P115>
2. Ahmad, S. F., Han, H., Alam, M. M., Rehmat, Mohd. K., Irshad, M., Arraño-Muñoz, M., & Ariza-Montes, A. (2023). Impact of Artificial Intelligence on Human Loss in Decision making, Laziness and Safety in Education. *Humanities and Social Sciences Communications*, 10(1), 1–14. <https://doi.org/10.1057/s41599-023-01787-8>
3. Aldboush, H. H. H., & Ferdous, M. (2023). Building Trust in Fintech: An Analysis of Ethical and Privacy Considerations in the Intersection of Big Data, AI, and Customer Trust. *International Journal of Financial Studies*, 11(3). <https://doi.org/10.3390/ijfs11030090>
4. Amnas, M. B., Selvam, M., & Parayitam, S. (2024). FinTech and Financial Inclusion: Exploring the Mediating Role of Digital Financial Literacy and the Moderating Influence of Perceived Regulatory Support. *Journal of Risk and Financial Management*, 17(3), 108. <https://doi.org/10.3390/jrfm17030108>
5. Anshari, M., Hamdan, M., Ahmad, N., Ali, E., & Haidi, H. (2022). COVID-19, Artificial intelligence, Ethical Challenges and Policy Implications. *AI & Society*, 38(2). <https://doi.org/10.1007/s00146-022-01471-6>
6. Bahoo, S., Cucculelli, M., Goga, X., & Mondolo, J. (2024). Artificial intelligence in Finance: a comprehensive review through bibliometric and content analysis. *Artificial Intelligence in Finance: A Comprehensive Review through Bibliometric and Content Analysis*, 4(2). <https://doi.org/10.1007/s43546-023-00618-x>
7. Barrow, J. M., & Khandhar, P. B. (2023, August 8). Deontology. National Library of Medicine; StatPearls Publishing. <https://www.ncbi.nlm.nih.gov/books/NBK459296/>
8. BBC. (2019, November 11). Apple's "sexist" credit card probed by regulator. BBC News. <https://www.bbc.com/news/business-50365609>
9. Bleher, H., & Braun, M. (2023). Reflections on Putting AI Ethics into Practice: How Three AI Ethics Approaches Conceptualize Theory and Practice. *Science and Engineering Ethics*, 29(3). <https://doi.org/10.1007/s11948-023-00443-3>
10. Boreiko, D., & Massarotti, F. (2020). How Risk Profiles of Investors Affect Robo-Advised Portfolios. *Frontiers in Artificial Intelligence*, 3. <https://doi.org/10.3389/frai.2020.00060>
11. Buolamwini, J., & Gebru, T. (2018). Gender Shades: Intersectional Accuracy Disparities in Commercial Gender Classification. *Proceedings of Machine Learning Research*, 81(1), 1–15. <https://proceedings.mlr.press/v81/buolamwini18a/buolamwini18a.pdf>
12. Card, D., & Smith, N. A. (2020). On Consequentialism and Fairness. *Frontiers in Artificial Intelligence*, 3(1), 1–11. <https://doi.org/10.3389/frai.2020.00034>

13. Černevičienė, J., & Kabašinskas, A. (2024). Explainable artificial intelligence (XAI) in finance: a systematic literature review. *Artificial Intelligence Review*, 57(8). <https://doi.org/10.1007/s10462-024-10854-8>
14. Chaddha, R., & Agrawal, G. K. (2023). Ethics and morality. *Indian Journal of Orthopaedics*, 57(11), 1707–1713. <https://doi.org/10.1007/s43465-023-01004-3>
15. Cheong, B. C. (2024). Transparency and accountability in AI systems: safeguarding wellbeing in the age of algorithmic decision-making. *Frontiers in Human Dynamics*, 6. <https://doi.org/10.3389/fhumd.2024.1421273>
16. Cristina, A., Gomes, M., & Rigobon, R. (2023). Algorithmic discrimination in the credit domain: what do we know about it? *AI & Society*, 39, 2059–2098. <https://doi.org/10.1007/s00146-023-01676-3>
17. D'Alessandro, W. (2024). Deontology and safe artificial intelligence. *Philosophical Studies*. <https://doi.org/10.1007/s11098-024-02174-y>
18. de Castro Vieira, J. R., Barboza, F., Cajueiro, D., & Kimura, H. (2025). Towards Fair AI: Mitigating Bias in Credit Decisions—A Systematic Literature Review. *Journal of Risk and Financial Management*, 18(5), 228. <https://doi.org/10.3390/jrfm18050228>
19. Dempsey, R. P., Eskander, E. E., & Dubljević, V. (2023). Ethical decision-making in law enforcement: A scoping review. *Psych*, 5(2), 576–601. <https://doi.org/10.3390/psych5020037>
20. Ebrahimi, S., Abdelhalim, E., Hassanein, K., & Head, M. (2024). Reducing the incidence of biased algorithmic decisions through feature importance transparency: an empirical study. *European Journal of Information Systems*, 34(4), 636–664. <https://doi.org/10.1080/0960085x.2024.2395531>
21. Farisco, M., Evers, K., & Salles, A. (2020). Towards Establishing Criteria for the Ethical Analysis of Artificial Intelligence. *Science and Engineering Ethics*, 26(5), 2413–2425. <https://doi.org/10.1007/s11948-020-00238-w>
22. Floridi, L., & Cows, J. (2019). A unified framework of five principles for AI in society. *Harvard Data Science Review*, 1(1). <https://doi.org/10.1162/99608f92.8cd550d1>
23. Floridi, L., Cows, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., Luetge, C., Madelin, R., Pagallo, U., Rossi, F., Schafer, B., Valcke, P., & Vayena, E. (2018). AI4People—An Ethical Framework for a Good AI Society: Opportunities, Risks, Principles, and Recommendations. *Minds and Machines*, 28(4), 689–707. <https://doi.org/10.1007/s11023-018-9482-5>
24. Fuster, A., Goldsmith-Pinkham, P., Ramadorai, T., & Walther, A. (2021). Predictably Unequal? The Effects of Machine Learning on Credit Markets. *The Journal of Finance*, 77(1), 5–47. <https://doi.org/10.1111/jofi.13090>
25. Giarmoleo, F. V., Ferrero, I., Rocchi, M., & Pellegrini, M. (2024). What ethics can say on artificial intelligence: Insights from a systematic literature review. *Business and Society Review*, 129(2), 258–292. <https://doi.org/10.1111/basr.12336>
26. Gladstone, J., & Hundtofte, C. S. (2023). A lack of financial planning predicts increased mortality risk: Evidence from cohort studies in the United Kingdom and United States. *PLoS ONE*, 18(9). <https://doi.org/10.1371/journal.pone.0290506>
27. Gorry, A., & Leganza, J. M. (2024). How do life events affect retirement timing? Evidence from expectations data. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.4815140>
28. Griffin, T. A., Green, B. P., & Welie, J. V. M. (2024). The ethical wisdom of AI developers. *AI and Ethics*, 5, 1087–1097. <https://doi.org/10.1007/s43681-024-00458-x>
29. Hagendorff, T. (2022). A Virtue-Based Framework to Support Putting AI Ethics into Practice. *Philosophy & Technology*, 35(3). <https://doi.org/10.1007/s13347-022-00553-z>
30. Hanna, M., Pantanowitz, L., Jackson, B., Palmer, O., Visweswaran, S., Pantanowitz, J., Deebajah, M., & Rashidi, H. (2024). Ethical and Bias Considerations in Artificial intelligence/machine Learning. *Modern Pathology*, 38(3), 1–13. ScienceDirect. <https://doi.org/10.1016/j.modpat.2024.100686>
31. Heidari, H., Loi, M., Gummadi, K. P., & Krause, A. (2019). A Moral Framework for Understanding Fair ML through Economic Models of Equality of Opportunity. *Proceedings of the Conference on Fairness, Accountability, and Transparency*, 181–190. <https://doi.org/10.1145/3287560.3287584>
32. Hermosilla, P., Díaz, M., Berríos, S., & Allende-Cid, H. (2025). Use of Explainable Artificial Intelligence for Analyzing and Explaining Intrusion Detection Systems. *Computers*, 14(5). <https://doi.org/10.3390/computers14050160>

33. Jedličková, A. (2024). Ethical approaches in designing autonomous and intelligent systems: a comprehensive survey towards responsible development. *AI & Society*, 40, 2703–2716. <https://doi.org/10.1007/s00146-024-02040-9>
34. Jia, T., Wang, C., Tian, Z., Wang, B., & Tian, F. (2022). Design of Digital and Intelligent Financial Decision Support System Based on Artificial Intelligence. *Computational Intelligence and Neuroscience*, 2022, 1–7. <https://doi.org/10.1155/2022/1962937>
35. Jobin, A., Ienca, M., & Vayena, E. (2019). The Global Landscape of AI Ethics Guidelines. *Nature Machine Intelligence*, 1(9), 389–399. <https://doi.org/10.1038/s42256-019-0088-2>
36. Kaas, M. H. L. (2024). The perfect technological storm: artificial intelligence and moral complacency. *Ethics and Information Technology*, 26(3). <https://doi.org/10.1007/s10676-024-09788-0>
37. Kazim, E., & Koshiyama, A. S. (2021). A High-level Overview of AI Ethics. *Patterns*, 2(9), 1–12. <https://doi.org/10.1016/j.patter.2021.100314>
38. Klingbeil, A., Grützner, C., & Schreck, P. (2024). Trust and Reliance on AI — an Experimental Study on the Extent and Costs of Overreliance on AI. *Computers in Human Behavior*, 160, 108352. <https://doi.org/10.1016/j.chb.2024.108352>
39. Kowald, D., Scher, S., Pammer-Schindler, V., Müllner, P., Waxnegger, K., Demelius, L., Fessler, A., Toller, M., Mendoza, G., Šimić, I., Sabol, V., Trügler, A., Veas, E., Kern, R., Nad, T., & Kopeinik, S. (2024). Establishing and evaluating trustworthy AI: overview and research challenges. *Frontiers in Big Data*, 7. <https://doi.org/10.3389/fdata.2024.1467222>
40. Lam, J. W. (2016, April 4). Robo-Advisors: A Portfolio Management Perspective. Yale Department of Economics. https://economics.yale.edu/sites/default/files/2023-01/Jonathan_Lam_Senior%20Essay%20Revised.pdf
41. Li, C., Wang, H., Jiang, S., & Gu, B. (2024). The Effect of AI-Enabled Credit Scoring on Financial Inclusion: Evidence from an Underserved Population of over One Million. *AIS Electronic Library*. <https://aisel.aisnet.org/misq/vol48/iss4/25/>
42. Liu, Z., & Liang, H. (2025). Are credit scores gender-neutral? Evidence of mis-calibration from alternative and traditional borrowing data. *Journal of Behavioral and Experimental Finance*, 47, 101081. <https://doi.org/10.1016/j.jbef.2025.101081>
43. Machado, J., Sousa, R., Peixoto, H., & António Abelha. (2024). Ethical Decision-Making in Artificial Intelligence: A Logic Programming Approach. *AI*, 5(4), 2707–2724. <https://doi.org/10.3390/ai5040130>
44. Madaan, H. (2025, February 14). XAI: Bringing Transparency And Trust To Algorithmic Decisions. *Forbes*. <https://www.forbes.com/councils/forbestechcouncil/2025/02/14/the-rise-of-explainable-ai-bringing-transparency-and-trust-to-algorithmic-decisions/>
45. Maier, T., Menold, J., & McComb, C. (2022). The Relationship Between Performance and Trust in AI in E-Finance. *Frontiers in Artificial Intelligence*, 5. <https://doi.org/10.3389/frai.2022.891529>
46. Mehrabi, N., Morstatter, F., Saxena, N., Lerman, K., & Galstyan, A. (2019). A Survey on Bias and Fairness in Machine Learning. *ArXiv:1908.09635 [Cs]*. <https://arxiv.org/abs/1908.09635>
47. Morley, J., Floridi, L., Kinsey, L., & Elhalal, A. (2021). From what to how: An overview of AI ethics tools, methods and research to translate principles into practices. *AI and Society*, 36, 59–72.
48. Najem, R., Bahnasse, A., Amr, M. F., & Talea, M. (2025). Advanced AI and big data techniques in E-finance: a comprehensive survey. *Discover Artificial Intelligence*, 5(1). <https://doi.org/10.1007/s44163-025-00365-y>
49. Nallakuruppan, M. K., Chaturvedi, H., Grover, V., Balusamy, B., Jaraut, P., Bahadur, J., Meena, V. P., & Hameed, I. A. (2024). Credit Risk Assessment and Financial Decision Support Using Explainable Artificial Intelligence. *Risks*, 12(10). <https://doi.org/10.3390/risks12100164>
50. Nwafor, C. N., Nwafor, O., & Brahma, S. (2024). Enhancing transparency and fairness in automated credit decisions: an explainable novel hybrid machine learning approach. *Scientific Reports*, 14(1). <https://doi.org/10.1038/s41598-024-75026-8>
51. Raji, A. A. H., Alabdoon, A. H. F., & Almagtome, A. (2024). AI in Credit Scoring and Risk Assessment: Enhancing Lending Practices and Financial Inclusion. *International Conference on Knowledge Engineering and Communication Systems*, 15, 1–7. <https://doi.org/10.1109/ickecs61492.2024.10616493>
52. Raji, I. D., Smart, A., White, R. N., Mitchell, M., Gebru, T., Hutchinson, B., Smith-Loud, J., Theron, D., & Barnes, P. (2020). Closing the AI Accountability Gap: Defining an End-to-End Framework for Internal Algorithmic Auditing. *ArXiv:2001.00973 [Cs]*. <https://doi.org/10.48550/arXiv.2001.00973>

53. Ridzuan, N. N., Masri, M., Anshari, M., Fitriyani, N. L., & Syafrudin, M. (2024). AI in the Financial Sector: The Line between Innovation, Regulation and Ethical Responsibility. *Information*, 15(8), 432. <https://doi.org/10.3390/info15080432>
54. Saarela, M., & Podgorelec, V. (2024). Recent Applications of Explainable AI (XAI): A Systematic Literature Review. *Applied Sciences*, 14(19). <https://doi.org/10.3390/app14198884>
55. Sahu, M. K. (2024). AI-Based Robo-Advisors: Transforming Wealth Management and Investment Advisory Services. *Journal of AI-Assisted Scientific Discovery*, 4(1), 379–411. <https://www.scienceadpress.com/index.php/jaasd/article/view/142>
56. Salloch, S., & Eriksen, A. (2024). What Are Humans Doing in the Loop? Co-Reasoning and Practical Judgment When Using Machine Learning-Driven Decision Aids. *American Journal of Bioethics*, 24(9), 1–12. <https://doi.org/10.1080/15265161.2024.2353800>
57. Santoni de Sio, F., & van den Hoven, J. (2018). Meaningful Human Control over Autonomous Systems: A Philosophical Account. *Frontiers in Robotics and AI*, 5(15). <https://doi.org/10.3389/frobt.2018.00015>
58. Spiekermann, S. (2017). IEEE P7000—The First Global Standard Process for Addressing Ethical Concerns in System Design. *Proceedings*, 1(3), 159. <https://doi.org/10.3390/is4si-2017-04084>
59. Svetlova, E. (2022). AI ethics and systemic risks in finance. *AI and Ethics*, 2(4), 713–725. <https://doi.org/10.1007/s43681-021-00129-1>
60. Szadeczky, T., & Bederna, Z. (2025). Risk, regulation, and governance: evaluating artificial intelligence across diverse application scenarios. *Security Journal*, 38(1). <https://doi.org/10.1057/s41284-025-00495-z>
61. Tao, R., Su, C.-W., Xiao, Y., Dai, K., & Khalid, F. (2021). Robo advisors, algorithmic trading and investment management: Wonders of fourth industrial revolution in financial markets. *Technological Forecasting and Social Change*, 163, 120421. <https://doi.org/10.1016/j.techfore.2020.120421>
62. Thein, H. H., Grosman, A., Sosnovskikh, S., & Klarin, A. (2024). Should we stay or should we exit? Dilemmas faced by multinationals under sanctioned regimes. *Journal of World Business*, 59(6), 101585. <https://doi.org/10.1016/j.jwb.2024.101585>
63. Ukanwa, K. (2024). Algorithmic Bias: Social Science Research Integration Through The 3-D Dependable AI Framework. *Current Opinion in Psychology*, 58, 101836. <https://doi.org/10.1016/j.copsyc.2024.101836>
64. Vuković, D. B., Dekpo-Adza, S., & Matović, S. (2025). AI integration in financial services: a systematic review of trends and regulatory challenges. *Humanities and Social Sciences Communications*, 12(1). <https://doi.org/10.1057/s41599-025-04850-8>
65. Wachter, S., Mittelstadt, B., & Floridi, L. (2017). Why a Right to Explanation of Automated Decision-Making Does Not Exist in the General Data Protection Regulation. *International Data Privacy Law*, 7(2), 76–99. <https://doi.org/10.1093/idpl/ix005>
66. Wachter, S., Mittelstadt, B., & Russell, C. (2018). Counterfactual Explanations without Opening the Black Box: Automated Decisions and the GDPR. *Harvard Journal of Law & Technology*. <https://arxiv.org/abs/1711.00399>
67. Walton, P. (2018). Artificial Intelligence and the Limitations of Information. *Information*, 9(12), 332. <https://doi.org/10.3390/info9120332>
68. Wang, X., Wu, Y. C., Ji, X., & Fu, H. (2024). Algorithmic discrimination: examining its types and regulatory measures with emphasis on US legal practices. *Frontiers in Artificial Intelligence*, 7, 1320277. <https://doi.org/10.3389/frai.2024.1320277>
69. Yordanova, K., & Bertels, N. (2023). Regulating AI: Challenges and the Way Forward Through Regulatory Sandboxes. *Law, Governance and Technology Series*, 58, 441–456. https://doi.org/10.1007/978-3-031-41264-6_23
70. Zhu, H., Vigren, O., & Söderberg, I.-L. (2024). Implementing artificial intelligence empowered financial advisory services: A literature review and critical research agenda. *Journal of Business Research*, 174. <https://doi.org/10.1016/j.jbusres.2023.114494>