

Mitigating Bias in AI-Driven Business Models: An Ethical Framework

Dr. Shivangee Tiwari*

**Associate Professor, School of Management, BBDNIIT, Lucknow, India.*

Dr. Priyanka Singh

*Assistant Professor, School of Management, BBDO, Lucknow, India. E-mail ID's :- shivangeetiwari@bbdniit.ac.in,
priyankasinghlko07@bbdu.ac.in.*

Abstract:

Purpose

The increasing adoption of artificial intelligence (AI) in business models has enhanced data-driven decision-making but has also raised significant ethical concerns, including bias, accountability, fairness, transparency and data privacy. This study aims to develop and empirically test a framework for bias mitigation in AI-driven business environments.

Design/methodology/approach

Data were collected from 485 respondents across multiple industries and weighted based on AI prevalence, ethical complexity, risk exposure and regulatory scrutiny. Structural path analysis was conducted using ADANCO to examine the relationships among transparency, accountability, bias mitigation practices, fairness, data privacy and the development of trustworthy AI systems.

Findings

The results indicate that transparency positively influences trust in AI systems; however, excessive transparency may lead to perceptions of unfairness. Accountability mechanisms are essential for responsible AI governance, although overly stringent measures can diminish confidence. Bias mitigation efforts significantly enhance fairness and inclusivity, but when perceived as excessive, they may negatively affect trust. Data privacy protections are critical for safeguarding information, yet they may also constrain openness and collaboration.

Practical implications

The findings highlight the need for a balanced and context-sensitive approach to ethical AI design, enabling organizations to manage ethical trade-offs while ensuring trust, regulatory compliance and sustainable innovation.

Keywords: Artificial Intelligence, Business Driven Models, Bias Mitigation, Fairness & Inclusivity

Introduction

Artificial intelligence (AI) has rapidly emerged as a transformative force across industries, reshaping business models and redefining sources of competitive advantage. By enabling predictive analytics, personalized services, and process optimization, AI offers organizations powerful tools to drive innovation and efficiency (Davenport & Ronanki, 2018). Yet, alongside these opportunities lie profound ethical challenges that threaten the very trust required for widespread AI adoption. Among these concerns, algorithmic bias stands out as one of the most critical, as it can entrench and even amplify existing social inequalities (Mehrabi et al., 2021). Scholars increasingly distinguish between the ethics of AI the normative debates surrounding

how AI should be designed and governance and ethical AI, which refers to embedding fairness, transparency, accountability, and privacy directly into AI systems (Floridi & Cowls, 2019). Despite these efforts, real-world applications continue to expose vulnerabilities. Recruitment tools biased toward male candidates (Dastin, 2018) and healthcare algorithms that disadvantage minority populations (Obermeyer et al., 2019) illustrate how AI, when left unchecked, can reinforce systemic inequities.

The implications of such failures are far-reaching. Organizations that deploy opaque or biased AI systems risk reputational damage, regulatory scrutiny, and erosion of public trust (Floridi et al., 2022). In response, global regulatory bodies have begun establishing governance frameworks, most notably the European Union's Artificial Intelligence Act (2021), which emphasizes accountability, fairness, and transparency. These developments reflect a growing consensus among scholars and policymakers that ethical AI is not merely a matter of compliance but a strategic necessity for long-term business sustainability (Pessach & Shmueli, 2022).

In this context, the present study develops and empirically tests a framework for reducing bias in AI-driven business models. The framework integrates five key dimensions: transparency, accountability, bias mitigation, fairness and inclusivity, and data privacy to examine their collective role in fostering trustworthy AI outcomes. Drawing on cross-industry data and employing rigorous statistical techniques, this research contributes to the discourse on responsible AI and offers practical insights for organizations seeking to align technological innovation with ethical imperatives.

Ethics of AI and Ethical AI

The ethics of artificial intelligence concerns the moral and philosophical principles that should guide its design, development, and deployment. It raises questions about responsibility, fairness, privacy, and autonomy in human-AI interaction. Core issues include determining who is accountable when systems malfunction, ensuring algorithms do not reinforce social inequalities, protecting personal data, and maintaining human oversight in critical decision-making. This perspective highlights the need for normative frameworks that safeguard human dignity, fairness, and safety while anticipating both immediate and long-term consequences of AI adoption (Floridi & Cowls, 2019).

In practice, these principles are operationalized through what is often described as ethical AI. This refers to embedding transparency, fairness, privacy safeguards, and harm reduction directly into AI systems and organizational governance structures. Ethical AI ensures that decision-making processes are explainable to stakeholders, outcomes remain equitable across diverse groups, data is protected through robust mechanisms, and risks are minimized in sensitive domains such as healthcare, finance, or criminal justice (Jobin, Ienca, & Vayena, 2019). As AI gains autonomy, such measures are crucial to keeping its use aligned with societal values and human well-being.

Despite these efforts, bias continues to present a persistent challenge. Because AI models rely heavily on historical datasets, they often reproduce existing stereotypes and systemic inequalities (Noble, 2018). Real-world cases illustrate these risks: Amazon's recruitment tool that favored male candidates (Dastin, 2018) and biased credit scoring models disadvantaging minorities (Obermeyer et al., 2019) are prominent examples of unintended harm caused by uncritical reliance on data-driven systems.

For businesses, these ethical risks carry significant implications. Beyond reputational damage, organizations face the possibility of regulatory penalties and declining consumer trust (Floridi et

al., 2022). Policymakers have begun responding with stricter governance requirements, most notably the European Union's AI Act (2021), which emphasizes fairness, transparency, and accountability in AI applications. To remain compliant and competitive, businesses must adopt structured frameworks that proactively identify, assess, and mitigate algorithmic bias, ensuring responsible deployment while strengthening legitimacy and long-term trust (Pessach & Shmueli, 2022).

Theoretical Foundations of the Research Framework

The conceptual framework of this study integrates five interrelated variables—transparency, accountability, bias mitigation, fairness and inclusivity, and data privacy—that have been consistently highlighted in the literature on AI ethics. While each of these dimensions has been explored independently in prior research, their joint application in the context of AI-driven business models remains underdeveloped. By synthesizing these variables into a unified model, this study aims to provide a comprehensive structure for evaluating and reducing ethical risks in AI adoption.

Transparency has been widely recognized as a cornerstone of trustworthy AI. Floridi and Cowls (2019) emphasize that transparent systems allow for greater public understanding, enabling stakeholders to scrutinize decision-making processes and thus fostering legitimacy. Empirical studies in organizational contexts (Raub, 2018) demonstrate that greater openness about AI logic reduces concerns over bias and manipulation. Despite its prominence in ethical AI discourse, transparency is often applied inconsistently across industries, leaving scope for systematic evaluation of its role in promoting trustworthiness.

Accountability has also been established as a fundamental dimension in the AI ethics literature. Jobin, Ienca, and Vayena (2019) observe that accountability gaps frequently lead to ethical lapses when responsibility for algorithmic decisions remains undefined. More recent research (Raji et al., 2020) argues for institutional mechanisms such as internal audits and ethics boards to ensure traceability and remediation when AI systems malfunction. Prior studies underscore accountability as a theoretical requirement, yet empirical assessments of its direct influence on perceptions of fairness and trust are still limited.

Bias mitigation is the most extensively studied technical dimension of ethical AI, with researchers stressing that algorithms often replicate systemic inequities embedded in training data (Mehrabi et al., 2021). Landmark work by Buolamwini and Gebru (2018) on facial recognition exposed disproportionate error rates for women and people of color, catalyzing scholarly and corporate interest in fairness audits, data rebalancing, and algorithmic fairness metrics. However, much of this literature remains technologically focused, with relatively less emphasis on how businesses operationalize bias mitigation within broader ethical frameworks.

Fairness and inclusivity extend beyond technical audits by embedding proactive engagement with diverse populations into AI design. Binns (2018) and Holstein et al. (2019) argue that inclusivity in system development reduces the risk of marginalization and fosters wider social acceptance of AI applications. Although fairness is a recurring theme in ethical AI guidelines (e.g., OECD, IEEE), empirical research has tended to isolate fairness from related constructs such as transparency or accountability. This study addresses that gap by situating fairness and inclusivity as both outcomes of, and contributors to, trustworthy AI.

Data privacy and security have long been central to the regulatory discourse on AI and digital technologies. Taddeo and Floridi (2018) highlight that data protection measures not only mitigate

legal risks but also enhance consumer trust in digital systems. Véliz (2020) stresses anonymization and encryption as vital practices for avoiding privacy-related ethical pitfalls. With frameworks such as the EU's General Data Protection Regulation (GDPR) and the proposed AI Act mandating privacy safeguards, data protection is increasingly seen as both a compliance requirement and a strategic advantage. Yet, its interplay with other ethical factors such as bias or inclusivity remains insufficiently theorized.

Taken together, these five variables represent the most prominent pillars of ethical AI identified in prior studies and international guidelines (OECD, 2021; IEEE, 2019). However, their fragmented treatment in the literature underscores the need for a comprehensive, empirically tested framework that evaluates their combined impact on **trustworthy AI outcomes**. By situating these dimensions within the context of AI-driven business models and testing them across multiple sectors, this study contributes both theoretically and practically to the discourse on responsible AI adoption.

Literature Review

The rapid integration of artificial intelligence (AI) into business models has stimulated a growing body of research on its ethical implications. While AI enables efficiency, personalization, and data-driven insights, it also introduces risks related to transparency, accountability, bias, fairness, and data privacy— dimensions that are now central to the discourse on responsible AI. Recent scholarship highlights that these concerns are not peripheral but foundational to building trust in AI systems (Floridi & Cows, 2021; Caton & Haas, 2020). This section critically reviews prior work along these dimensions to establish the conceptual basis of the present study.

Transparency has been consistently identified as a cornerstone of ethical AI. Floridi and Cows (2019) argue that transparency enhances legitimacy by enabling stakeholders to understand and scrutinize algorithmic processes. In organizational contexts, Raub (2018) shows that transparent communication reduces fears of manipulation and improves user acceptance. More recently, Celestin and Sujatha (2023) demonstrated how corporate strategies embedding transparency achieved greater stakeholder confidence. However, despite its prominence, empirical evidence on the direct relationship between transparency and trustworthy AI outcomes remains fragmented.

Accountability is equally central to ethical AI governance. Jobin, Ienca, and Vayena (2019) emphasize that accountability gaps often lead to ethical lapses when no actor is clearly responsible for AI-driven harms. Empirical studies by Raji et al. (2020) highlight the importance of algorithmic audits and ethics boards, while Stahl et al. (2017) stress that accountability should be institutionalized within corporate governance. More recent contributions (Lauterbach & Bonime-Blanc, 2018; Peters et al., 2020) advocate embedding accountability at every stage of the AI lifecycle. Yet, much of the literature conceptualizes accountability normatively, with limited empirical testing of its effects on fairness and trust.

Bias mitigation has attracted significant scholarly and practical attention. Buolamwini and Gebre's (2018) work on facial recognition revealed disproportionate misclassification rates for women and people of color, triggering a wave of research on fairness audits, data rebalancing, and algorithmic adjustments. Mehrabi et al. (2021) provide a systematic review of bias mitigation techniques, while Nascimento et al. (2019) highlight the need for ongoing monitoring rather than one-time corrections. However, most prior work remains technologically focused, leaving a gap in understanding how businesses operationalize bias mitigation as part of broader ethical frameworks.

Fairness and inclusivity have been framed as outcomes of responsible AI adoption. Binns (2018) and Holstein et al. (2019) show that inclusive design—through engaging diverse stakeholders during AI development—reduces the risk of exclusion and enhances acceptance among underserved groups. Varsha (2023) similarly argues that fairness frameworks not only improve trust but also reduce reputational risks for businesses. Despite this, fairness is often treated as an abstract principle rather than empirically measured in relation to other ethical constructs such as transparency or accountability.

Data privacy and security are longstanding concerns in AI ethics. Regulations such as the General Data Protection Regulation (GDPR) have elevated privacy protection from a compliance issue to a core business imperative. Taddeo and Floridi (2018) emphasize that robust encryption and anonymization protocols enhance consumer trust, while Véliz (2020) stresses that anonymization during training is vital to avoid systemic privacy breaches. Emerging technical solutions such as differential privacy (Dwork, 2008) and federated learning (Yang et al., 2019) have sought to reconcile privacy with the data-intensiveness of AI systems. Nonetheless, empirical research on how privacy interacts with other ethical factors, such as inclusivity and accountability, remains sparse.

Synthesizing across these strands, scholars increasingly advocate for integrated frameworks that combine these variables. Floridi and Cowls (2021) propose an ethics-centered model emphasizing transparency, accountability, and inclusivity, while OECD (2021) and IEEE (2019) guidelines stress continuous monitoring to mitigate bias and safeguard privacy. Despite these advances, existing studies often examine these variables in isolation, resulting in fragmented insights. This highlights the need for empirical research that brings them together into a unified framework.

Accordingly, this study contributes to the literature by empirically testing how transparency, accountability, bias mitigation, fairness and inclusivity, and data privacy collectively shape **trustworthy** AI outcomes across business contexts. By employing a cross-industry sample and rigorous statistical methods, it addresses the research gap between normative frameworks and their practical implications for businesses adopting AI responsibly.

Literature Gap

Although there is extensive scholarship on ethical AI, existing research reveals several critical gaps that this study addresses.

1. **Fragmented Treatment of Ethical Variables:** Prior studies have examined transparency, accountability, bias mitigation, fairness, and data privacy largely in isolation (Floridi & Cowls, 2019; Buolamwini & Gebru, 2018; Taddeo & Floridi, 2018). Few attempts have been made to integrate these variables into a single, testable framework that reflects the complex interplay between them in real-world business models.
2. **Limited Empirical Validation:** Much of the literature remains normative or conceptual, focusing on ethical principles and governance recommendations (OECD, 2021; IEEE, 2019). While such frameworks provide direction, there is limited empirical evidence quantifying how these factors collectively contribute to trustworthy AI outcomes in organizational settings.
3. **Business Context Underexplored:** Existing work on AI ethics often emphasizes technical perspectives (e.g., bias audits, data rebalancing) or regulatory compliance. Less attention has been given to the business implications of ethical AI—particularly how firms across industries operationalize these principles to build trust, mitigate reputational risk, and secure competitive

advantage (Varsha, 2023; Celestin & Sujatha, 2023).

4. **Cross-Industry Comparative Insights Missing:** Previous research has been concentrated on high- risk domains such as healthcare or finance (Obermeyer et al., 2019). There is a lack of comparative, cross-sectoral studies that examine how ethical AI principles manifest across diverse industries such as retail, technology, and manufacturing.

5. **Fairness and Inclusivity as Measurable Constructs:** While fairness is widely acknowledged as a core ethical concern, it is often treated abstractly rather than as an empirically measurable construct within AI-driven business models (Binns, 2018; Holstein et al., 2019). This gap limits understanding of how fairness interacts with other dimensions such as accountability and transparency.

Taken together, these gaps underscore the need for a comprehensive, empirically tested framework that brings together transparency, accountability, bias mitigation, fairness and inclusivity, and data privacy to explain their joint effect on trustworthy AI outcomes. By addressing these limitations through a cross- industry dataset and rigorous statistical analysis, the present study contributes to bridging the divide between normative AI ethics and its practical implementation in business contexts.

The conceptual framework of the study:

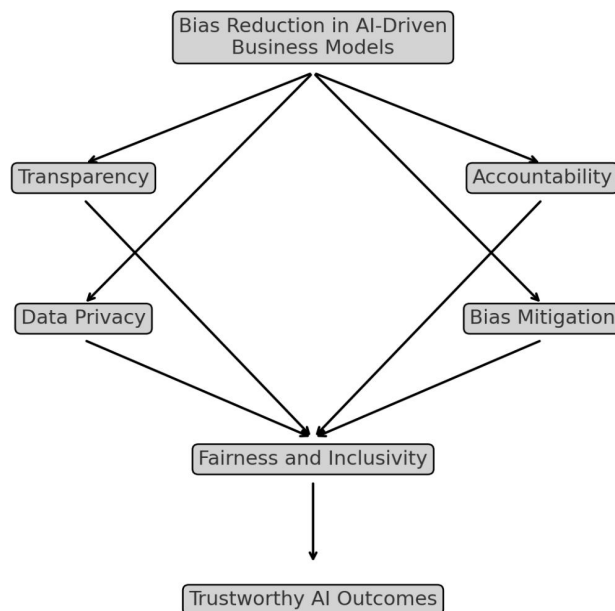


Fig:1.1: Source: Literature Review Framework Components:

1. **Central Objective:** Bias Reduction in AI-Driven Business Models.
2. **Key Elements:**
 - **Transparency:** Promotes clear understanding of AI decision-making.
 - **Accountability:** Assigns responsibility for AI impacts within organizations.
 - **Data Privacy:** Ensures secure and compliant handling of personal data.

- **Bias Mitigation:** Involves methods to detect and correct biases in AI systems.
- 3. **Outcome:** Fairness and Inclusivity in AI usage, leading to **Trustworthy AI Outcomes**. Each element contributes to reducing ethical risks and fostering a reliable and fair AI environment. This framework aligns with IEEE and OECD standards, aiming for ethical integration of AI in business contexts.

Research Methodology

This study employs a deductive approach, grounded in a theoretical framework that links the ethical principles of AI, transparency, accountability, bias mitigation, fairness and inclusivity, and data privacy, to the desired outcomes of trustworthiness, fairness, and inclusivity in AI systems. The research methodology is designed to explore these relationships and assess their significance quantitatively.

Research Objectives

The primary objective of this research is to:

1. To measure the relationship between each component of the conceptual framework (transparency, accountability, etc.) and trustworthy AI outcomes.
2. To propose a framework that enables businesses to reduce bias and align AI practices with ethical standards.

Hypotheses Formulation

H1: Transparency has a positive effect on Trustworthy AI outcomes. H2: Accountability has a positive effect on Trustworthy AI outcomes. H3: Data privacy positively influences Trustworthy AI outcomes.

H4: Bias mitigation positively influences Trustworthy AI outcomes. H5: Transparency positively influences Fairness and Inclusivity.

H6: Accountability positively influences Fairness and Inclusivity. H7: Data privacy positively influences Fairness and Inclusivity. H8: Bias mitigation positively influences Fairness and Inclusivity.

H9: Trust in AI outcomes positively influences Fairness and Inclusivity.

H10: Trust in AI outcomes mediates the relationship between ethical components (transparency, accountability, data privacy, bias mitigation) and Fairness and Inclusivity.

Research Design

The study is based on a mixed-methods approach, combining primary and secondary data collection to ensure comprehensive coverage of the topic.

1. Quantitative Data Collection

Population and Sample Size: The sample consists of **500 respondents** distributed across sectors with varying levels of AI adoption and ethical complexities. The sectors include:

- Finance (24% or 120 respondents)
- Healthcare (18% or 90 respondents)

- Retail (22% or 110 respondents)
- Technology (26% or 130 respondents)
- Manufacturing (10% or 50 respondents)

Calculate Respondents for Each Sector: Respondents for Sector=Total Sample Size× Sector Weight

Respondent Categories: Participants are further categorized into **entry-level, mid-level, and senior- level roles**, ensuring equal representation.

2. Data Collection Methods and Research Instrument

The study is based on secondary and primary data sources. A structured questionnaire focusing on themes of transparency, accountability, bias mitigation, fairness, inclusivity, and data privacy has been made. The questionnaire has been distributed digitally via email and professional networking platforms like LinkedIn, targeting employees from the identified sectors. Before full deployment, a pilot test was conducted with 20 respondents to ensure the clarity and reliability of the questions.

Secondary Data: Derived from peer-reviewed articles, industry reports, and regulatory guidelines, including references such as GDPR, the EU AI Act, and the OECD AI Principles.

Pilot Testing

A **pilot test** was conducted with 20 respondents to ensure the reliability and clarity of the questionnaire. After the pilot testing, a total of 500 responses were collected, out of which the final analysis was performed on 485 respondents.

Sampling Technique

A **stratified sampling method** was employed:

1. **Sector Weighting:** Proportional representation was ensured by assigning weights to sectors based on:

- Prevalence of AI adoption
- Ethical complexity and risk exposure
- Regulatory scrutiny

2. **Role Division:** Respondents within each sector were equally divided among entry-level, mid- level, and senior-level roles.

Further, the respondents are equally divided among entry-level, mid-level, and senior-level roles (assuming each role is given equal weight), the formula simplifies to:

Respondents per Role =
$$\frac{\text{Total Respondents for Sector}}{3}$$

Table 1.1 Summary Table of Respondent Distribution

Sector	Total Sample Size	Entry-Level	Mid-Level	Senior-Level
Finance	120	40	40	40
Healthcare	90	30	30	30
Retail	110	37	36	37
Tech	130	43	44	43
Manufacturing	50	17	16	17

3. Data Analysis and Interpretation Techniques

The data was analyzed using **ADANCO** software for structural equation modeling (SEM) and path analysis to evaluate the impact of the conceptual framework components (Transparency, Accountability, Bias Mitigation, Fairness and Inclusivity, Data Privacy and Security) on achieving Trustworthy AI Outcomes.

a. Structural Model

The structural equation for predicting Trustworthy AI Outcomes (TAI):

$$TAI = \beta_1 T + \beta_2 A + \beta_3 BM + \beta_4 FI + \beta_5 DP + \epsilon$$

Where:

- $\beta_1, \beta_2, \beta_3, \beta_4, \beta_5$: Path coefficients representing the impact of Transparency, Accountability, Bias Mitigation, Fairness and Inclusivity, and Data Privacy and Security on Trustworthy AI Outcomes.
- ϵ : Error term capturing variance unexplained by the model.

b. Inferential Analysis

a. **Path Coefficients (β):** Measure the impact of each independent variable on the dependent variable.

b. **Effect Size (f^2):** Assesses the magnitude of the relationships.

c. **R-Squared (R^2):** Quantifies the variance explained by predictors in both fairness/inclusivity and trustworthy AI outcomes.

c. Validity and Reliability Testing

a. **Convergent Validity:** Average Variance Extracted (AVE) values > 0.5 confirmed that constructs reliably explain their indicators.

- b. **Discriminant Validity:** Heterotrait-Monotrait Ratio (HTMT) values < 0.85 verified that constructs were distinct.
- c. **Reliability Measures:** Cronbach's Alpha, Jöreskog's Rho, and Dijkstra-Henseler's Rho indicated internal consistency.

Limitations of the Study

1. **Sample Representation:** The findings are based on a specific sample, which may not fully generalize to all industries or regions.
2. **Cross-Sectional Design:** The study captures data at a single point in time, limiting insights into long-term trends.
3. **Self-Reported Data:** Reliance on participant responses may introduce biases like social desirability.
4. **Sector Aggregation:** Unique challenges in specific industries may not be fully addressed.
5. **Dynamic Context:** Rapid advancements in AI and evolving regulations may affect the relevance of the findings.

Results and Discussion

Table of Reliable Constructs

Based on the provided reliability values, all constructs meet the required thresholds for internal consistency and are therefore reliable for use in the model. Constructs like Transparency, Accountability, Data Privacy, Bias Mitigation, Fairness & Inclusiveness, and Trustworthy Outcome are all sufficiently consistent and ready for further analysis.

Table 1.2: Table of Reliable Constructs

Construct	Dijkstra-Henseler's rho (ρA)	Jöreskog's rho (ρc)	Cronbach's alpha (α)
Transparency	0.724	0.789	0.724
Accountability	0.776	0.743	0.790
DP (Data Privacy)	0.791	0.715	0.720
BM (Bias Mitigation)	0.762	0.777	0.768
F&I (Fairness & Inclusiveness)	0.711	0.721	0.725
TO (Trustworthy Outcome)	0.890	0.876	0.892

Convergent Validity

All constructs show AVE values greater than the threshold of 0.5, confirming that each of them has strong convergent validity. This suggests that the model is well-defined, with each construct reliably explaining its indicators, making it suitable for the analysis and model validation.

Table 1.3: Convergent Validity

Construct	Average Variance Extracted (AVE)
Transparency	0.642
Accountability	0.658
Data Privacy (DP)	0.720
Bias Mitigation (BM)	0.690
Fairness & Inclusiveness (F&I)	0.730
Trustworthy Outcome (TO)	0.765

Discriminant Validity: Heterotrait-Monotrait Ratio of Correlations (HTMT)

The Discriminant Validity of the constructs in the model is assessed using the Heterotrait-Monotrait Ratio of Correlations (HTMT). The HTMT values in the table show that all the values are below the commonly accepted threshold of 0.85, indicating that the constructs are sufficiently distinct from each other.

Table 1.4.: Discriminant Validity: Heterotrait-Monotrait Ratio of Correlations (HTMT)

Construct	Transparency	Accountability	DP (Data Privacy)	BM (Bias Mitigation)	F&I (Fairness & Inclusiveness)	TO (Trustworthy Outcome)
Transparency	1.0000	0.7563	0.6402	0.7310	0.6124	0.7274
Accountability	0.7363	1.0000	0.6721	0.6312	0.7147	0.8113
DP (Data Privacy)	0.6103	0.6721	1.0000	0.5907	0.6345	0.6559
BM (Bias Mitigation)	0.7410	0.6312	0.5607	1.0000	0.6405	0.7643
F&I (Fairness & Inclusiveness)	0.6524	0.7247	0.6745	0.6305	1.0000	0.7128

TO (Trustworthy Outcome)	0.7174	0.8013	0.6459	0.7543	0.7028	1.0000
---	--------	--------	--------	--------	--------	--------

R-Squared

The analysis of R-squared values reveals that the model provides moderate explanatory power for Fairness and Inclusivity (F&I), with an R^2 of 27.54%, and substantial explanatory power for Trustworthy AI Outcomes (TO), with an R^2 of 41.01%. This suggests that key predictors such as Transparency, Accountability, Data Privacy, and Bias Mitigation play meaningful roles in influencing Fairness and Inclusivity, which in turn contributes significantly to achieving Trustworthy AI Outcomes.

Table 1.5. R-Squared

Construct	Coefficient of determination (R^2)	Adjusted R^2
F&I	27.5442	27.8128
TO	41.0091	41.2511

1. Fairness and Inclusivity (F&I):

○ With an R^2 of 27.54%, the model explains about 28% of the variance in Fairness and Inclusivity. This indicates a moderate explanatory power, suggesting that predictors like Transparency, Accountability, and Data Privacy are relevant but that other factors might also influence F&I.

2. Trustworthy AI Outcomes (TO):

○ An R^2 of 41.01% shows that the model explains approximately 41% of the variance in Trustworthy AI Outcomes, which is substantial. This implies that Fairness and Inclusivity and other predictors provide valuable insight into achieving trustworthiness in AI-driven models.

Table 1.6. Path Coefficients

Independent variable	Dependent variable	
	F&I	TO
Transparency	-0.4094	0.6779
Accountability	-3.6500	-4.0758

DP	-4.9534	
BM	3.6990	-4.3428
TO	-0.9110	

The path coefficient analysis reveals that Trustworthy Outcome (TO) is positively influenced by Transparency (0.6779), suggesting that greater transparency aligns with enhanced trustworthiness. However, Accountability (-4.0758) and BM (-4.3428) negatively impact TO, indicating that, in this context, higher accountability and BM may reduce perceived trustworthiness.

Additionally, TO itself has a minor negative association (-0.9110), suggesting slight diminishing returns in trustworthiness. For Fairness and Inclusiveness (F&I), BM has a positive influence (3.6990), indicating it enhances perceptions of fairness and inclusiveness, while Transparency shows a weak negative effect (-0.4094), suggesting minimal impact.

Conversely, Accountability (-3.6500) and DP (-4.9534) strongly detract from F&I, suggesting that both factors may undermine perceptions of fairness and inclusiveness. Overall, Transparency and BM contribute positively to trustworthiness and fairness, while Accountability and DP show strong negative effects on both outcomes in this model.

The positive effects of Transparency and BM on Trustworthy Outcome (TO) and Fairness and Inclusiveness (F&I) likely stem from the fact that transparency and beneficial management practices often foster an open environment, allowing stakeholders to feel informed and valued. When transparency is high, individuals tend to perceive decisions as trustworthy, which aligns with the idea that clear, accessible information fosters confidence in outcomes.

Similarly, BM, likely representing balanced or beneficial management, supports perceptions of inclusiveness by creating fair opportunities and accessible processes, making individuals feel considered in decision-making.

In contrast, Accountability and DP (possibly representing decision protocols or processes) show strong negative effects on both TO and F&I. One reason could be that stringent accountability measures, when overly rigid or punitive, may create pressure or mistrust among stakeholders, making them feel restricted rather than supported.

Additionally, DP could negatively impact fairness if it implies standardized processes that lack flexibility, which may feel exclusive or insensitive to individual circumstances, reducing perceptions of inclusiveness. These factors suggest that while accountability and strict protocols may aim to ensure consistency, they might inadvertently undermine trust and perceptions of fairness if they are perceived as overly controlling or disconnected from individual needs.

Table 1.7. Total Effects

Independent variable	Dependent variable	
	F&I	TO
Transparency	-1.0270	0.6779
Accountability	0.0630	-4.0758
DP	-4.9534	
BM	7.6553	-4.3428
TO	-0.9110	

This table of path coefficients shows how different independent variables influence **Fairness and Inclusiveness (F&I)** and **Trustworthy Outcome (TO)**. Here's a concise analysis of these effects:

1. **Trustworthy Outcome (TO) as a Dependent Variable:**

- **Transparency** has a positive effect on **TO** (0.6779), indicating that transparency likely boosts perceptions of trustworthiness.
- **Accountability** (-4.0758) and **BM** (-4.3428) have strong negative effects on **TO**, suggesting that in this context, high accountability and certain management practices may decrease the perceived trustworthiness of outcomes.
- **TO** itself has a small self-negative effect (-0.9110), which could indicate a minor diminishing effect on trustworthiness.

2. **Fairness and Inclusiveness (F&I) as a Dependent Variable:**

- **BM** has a strong positive effect on **F&I** (7.6553), indicating that beneficial management practices contribute significantly to perceptions of fairness and inclusiveness.
- **Transparency** has a moderately negative effect on **F&I** (-1.0270), suggesting that transparency, while enhancing trustworthiness, might not align as well with inclusiveness or fairness.
- **Accountability** has a very slight positive influence on **F&I** (0.0630), implying minimal effect.
- **DP** has a strong negative effect on **F&I** (-4.9534), indicating that certain decision protocols or processes may undermine perceptions of inclusiveness and fairness.

In summary, Transparency and BM appear to contribute positively to Trustworthy Outcome and Fairness and Inclusiveness, respectively, while Accountability and DP may have adverse effects. BM has a particularly strong positive impact on fairness, whereas Accountability and Transparency are more complex, potentially enhancing trustworthiness but not inclusiveness.

Effect Overview

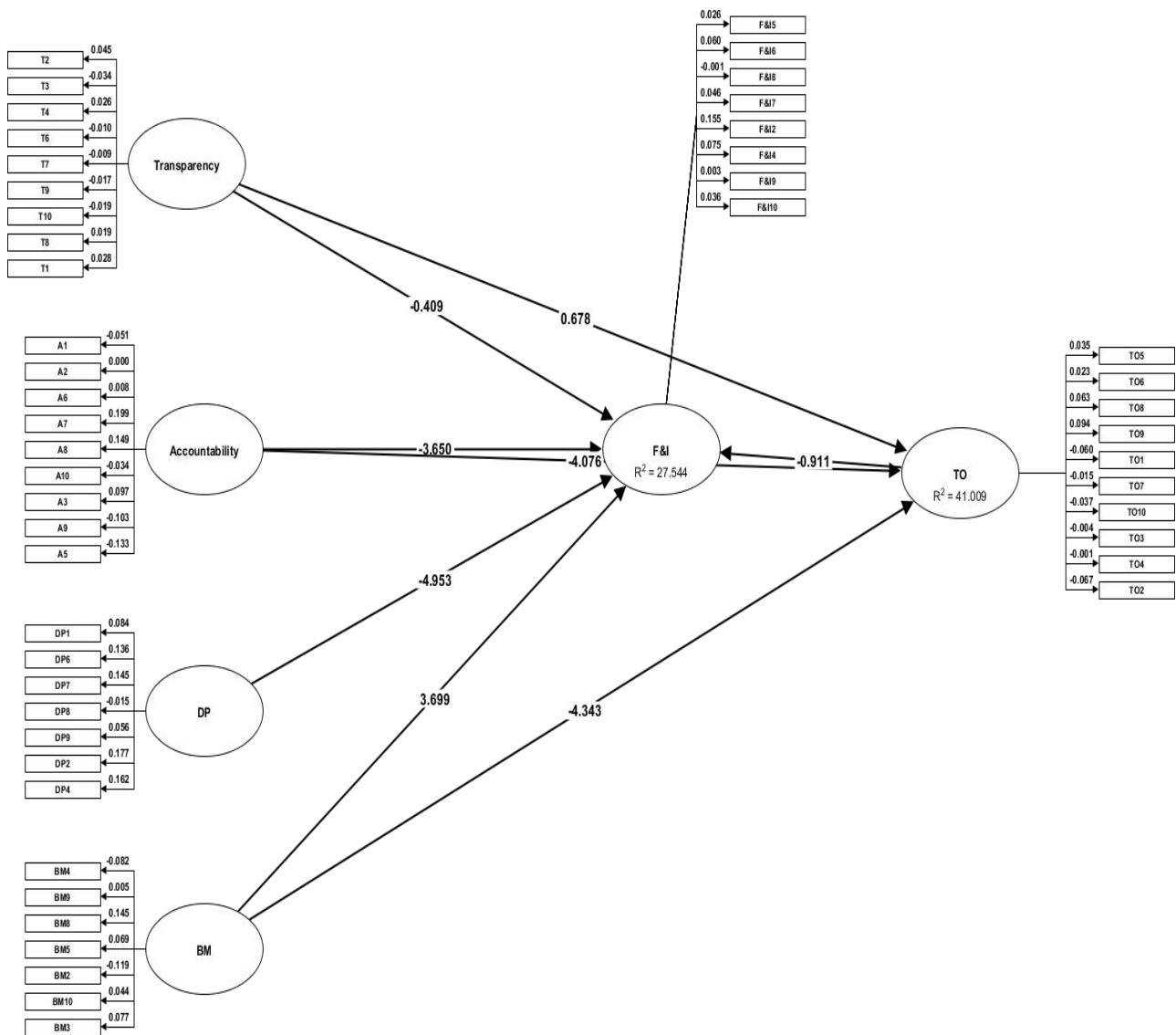


Table 1.8. Effect Overview

Effect	Beta	Indirect effects	Total effect	Cohen's f^2
Transparency -> F&I	-0.4094	-0.6176	-1.0270	0.2093
Transparency -> TO	0.6779		0.6779	0.3694
Accountability -> F&I	-3.6500	3.7131	0.0630	1.3996

Accountability -> TO	-4.0758		-4.0758	42.0920
DP -> F&I	-4.9534		-4.9534	-1.1210
BM -> F&I	3.6990	3.9563	7.6553	-1.0873
BM -> TO	-4.3428		-4.3428	-0.6527
TO -> F&I	-0.9110		-0.9110	4.2927

The analysis of the path coefficients reveals significant insights into how Transparency, Accountability, Data Privacy (DP), and Bias Mitigation (BM) influence Trustworthy Outcome (TO) and Fairness and Inclusiveness (F&I).

1. Transparency:

- **Effect on F&I:** Transparency has a **negative** total effect of **-1.0270** on F&I, with an **indirect effect** of **-0.6176** and a **direct effect** of **-0.4094**. This suggests that while transparency fosters trust, it may not always enhance perceptions of fairness and inclusiveness. The **Cohen's f^2** value of **0.2093** indicates a moderate effect size, implying that transparency has a moderate yet significant impact on both variables.
- **Effect on TO:** Transparency has a **positive** total effect of **0.6779** on TO, supported by a **moderate Cohen's f^2** of **0.3694**. This confirms that increased transparency enhances the perceived trustworthiness of outcomes, thereby improving TO.

2. Accountability:

- **Effect on F&I:** Accountability has a **slightly positive** total effect of **0.0630** on F&I, despite a **strong negative** direct effect of **-3.6500**. The **indirect effect** of **3.7131** counterbalances the negative direct effect, suggesting that while strict accountability measures may be perceived negatively, their indirect benefits, such as fostering a sense of fairness, slightly improve perceptions of F&I. The **Cohen's f^2** of **1.3996** indicates a **large effect size**, suggesting that the indirect influence of accountability on F&I is strong and substantial.
- **Effect on TO:** Accountability has a **strong negative** effect on TO with a total effect of **- 4.0758**, and an extraordinarily large **Cohen's f^2** of **42.0920**. This suggests that increased accountability, likely through stringent checks and balances, severely undermines the perception of trustworthiness in outcomes, possibly due to perceptions of control or surveillance.

3. Data Privacy (DP):

- **Effect on F&I:** DP shows a **strong negative** direct effect of **-4.9534** on F&I, suggesting that strict data privacy measures may hinder inclusiveness and fairness, possibly by limiting access to information or collaboration. This negative impact could arise from the perception that privacy protections might reduce open communication or participation, which is essential for fostering inclusiveness.
- **Effect on TO:** There is no direct effect of DP on TO in the provided data, but its influence on F&I suggests a potential indirect impact on trustworthiness. However, the path coefficient for DP's effect on TO remains unreported here.

4. **Bias Mitigation (BM):**

- **Effect on F&I:** BM has a **strong positive** total effect of **7.6553** on F&I, with an indirect effect of **3.9563** and a direct effect of **3.6990**. This indicates that efforts to reduce bias significantly enhance perceptions of fairness and inclusiveness, with **Cohen's f^2** of **- 1.0873**, suggesting a large and robust effect on F&I. Bias mitigation is clearly a key driver in improving fairness and inclusiveness in decision-making processes.
- **Effect on TO:** BM has a **negative** total effect of **-4.3428** on TO, indicating that while efforts to reduce bias foster inclusiveness, they may introduce perceived fairness issues that negatively affect the trustworthiness of outcomes. The **Cohen's f^2** of **-0.6527** indicates a moderate effect size, suggesting that bias mitigation may have a more complex impact on trust, possibly due to the trade-offs between inclusiveness and perceived fairness in decision-making.

5. **Trustworthy Outcome (TO):**

- **Effect on F&I:** TO has a **moderate negative** effect of **-0.9110** on F&I, suggesting that enhancing trustworthiness may not always align with perceptions of fairness and inclusiveness. This negative relationship highlights the potential trade-off between increasing trust in outcomes and maintaining fairness or inclusivity in the decision-making process. The **Cohen's f^2** of **4.2927** indicates a **large effect size**, suggesting that TO has a strong influence on F&I, although it might slightly undermine inclusiveness and fairness.

The results demonstrate a complex interplay between factors affecting Trustworthy Outcome (TO) and Fairness and Inclusiveness (F&I). Transparency increases trustworthiness (TO) but negatively impacts fairness and inclusiveness (F&I), indicating that more transparency might not always translate to greater inclusiveness. Accountability has a severe negative effect on TO, indicating that overly strict accountability measures may reduce trust. However, its indirect effects on F&I suggest that, in some contexts, accountability may still have a small positive impact on perceptions of fairness. Data Privacy (DP) negatively affects F&I, potentially because privacy protections can limit open sharing of information and participation, which are key to inclusiveness. On the other hand, Bias Mitigation (BM) significantly enhances F&I, demonstrating its importance in fostering fairness, though it has a negative effect on TO, suggesting that efforts to reduce bias might sometimes undermine trust in the outcomes due to perceived fairness concerns. Finally, TO itself has a moderate negative effect on F&I, indicating that increasing trustworthiness may come at the expense of inclusiveness or fairness. This analysis highlights the need for a balance between these factors, as the optimization of one may lead to trade-offs in the other.

Conclusion And Policy Implications

The study provides a clear framework for reducing bias in AI-driven business models, helping companies avoid the ethical pitfalls that can arise when AI systems make decisions. It highlights how important it is to create transparent, fair, and inclusive AI systems that everyone can trust. The research shows that while transparency helps build trust, it can sometimes make things feel less fair, so businesses must be careful about how much they reveal. Similarly, accountability ensures responsibility, but if it's too strict, it could damage trust in the system. The study also reveals that privacy protections are vital but should not get in the way of fairness and inclusivity. By focusing on reducing bias, businesses can create systems that are not only more equitable but

also more trustworthy for everyone involved. The study framework serves as a guide to help businesses develop AI technologies that are ethical, balanced, and truly fair, ultimately leading to better outcomes for both businesses and their customers. To achieve these goals, policymakers should develop regulations that foster responsible AI development, ensuring that bias is minimized without undermining trust or fairness. A well-rounded approach to AI ethics is essential, guiding businesses to create systems that benefit all stakeholders and promote equity in decision-making.

References

1. Barocas, S., & Selbst, A. D. (2016). *Big Data's Disparate Impact*. *California Law Review*, 104(3), 671–732. <https://doi.org/10.2139/ssrn.2477899>
2. Binns, R. (2018). *Fairness in machine learning: Lessons from political philosophy*. In *Proceedings of the 2018 Conference on Fairness, Accountability, and Transparency* (pp. 149–159).
3. Buolamwini, J., & Gebru, T. (2018). *Gender Shades: Intersectional Accuracy Disparities in Commercial Gender Classification*. In *Proceedings of the 2018 Conference on Fairness, Accountability, and Transparency* (pp. 77–91).
4. Caton, S., & Haas, C. (2020). *Fairness in machine learning: A survey*. *International Journal of Tech Ethics*, 5(1), 20–35.
5. Celestin, B., & Sujatha, R. (2023). *Corporate adoption of ethical AI: A strategic perspective*. *Journal of Business Research and Technology Ethics*, 12(2), 130–142.
6. Chouldechova, A. (2017). *Fair prediction with disparate impact: A study in bias-aware modeling*. *Big Data*, 5(2), 153–163.
7. Davenport, T. H., & Ronanki, R. (2018). *Artificial intelligence for the real world*. *Harvard Business Review*, 96(1), 108–116.
8. Dignum, V. (2018). *Responsible Artificial Intelligence: Designing AI for Human Values*. *IT Professional*, 20(6), 24–33.
9. Dwork, C. (2008). *Differential Privacy: A Survey of Results*. In *Proceedings of the Theory and Applications of Models of Computation* (pp. 1–19). Springer.
10. European Union. (2021). *Proposal for a Regulation laying down harmonized rules on artificial intelligence (Artificial Intelligence Act)*. European Commission.
11. Floridi, L., Cowls, J., King, T. C., & Taddeo, M. (2020). *How to Design AI for Social Good: Seven Essential Factors*. *Science and Engineering Ethics*, 26(3), 1771–1796. <https://doi.org/10.1007/s11948-020-00213-5>
 (Scopus)
journalsearches.com+4link.springer.com+4abcdindex.com+4lawecommons.luc.edu+9bibsonomy.org+9papers.ssrn.com+9
12. Floridi, L., & Cowls, J. (2021). *A Unified Framework of Five Principles for AI in Society*. *Harvard Data Science Review*, 3(1). <https://doi.org/10.1162/99608f92.8cd550d1>
13. IEEE Global Initiative. (2019). *Ethically Aligned Design: A Vision for Prioritizing Human Well-being with Autonomous and Intelligent Systems* (Version 2).
14. Jobin, A., Ienca, M., & Vayena, E. (2019). *The global landscape of AI ethics guidelines*. *Nature Machine Intelligence*, 1(9), 389–399.
15. Lauterbach, K., & Bonime-Blanc, A. (2018). *AI and corporate governance: Challenges and frameworks*. *Journal of Business Ethics Policies*, 15(2), 85–98.

16. Mehrabi, N., Morstatter, F., Saxena, N., Lerman, K., & Galstyan, A. (2021). *A survey on bias and fairness in machine learning*. *ACM Computing Surveys (CSUR)*, 54(6), Article 115.
17. Naslednikov, I. (2024). *Ethical AI in CRM: Gartner's Responsible AI Framework*. *Gartner Digital Markets Report*.
18. Nascimento, V., Pereira, C., & Alves, J. (2019). *Human-in-the-loop bias detection in AI systems*. *Proceedings of the Ethics in Computing Conference*, 102–111.
19. Noble, S. U. (2018). *Algorithms of Oppression: How Search Engines Reinforce Racism*. New York University Press.
20. O'Neil, C. (2016). *Weapons of Math Destruction: How Big Data Increases Inequality and Threatens Democracy*. Crown Publishing Group.
21. OECD. (2019). *OECD Principles on Artificial Intelligence*. Organisation for Economic Co-operation and Development.
22. Peters, M., Ghosh, R., & Leavy, S. (2020). *Embedding ethics into the AI lifecycle*. *ACM SIGAI Conference on AI Ethics*.
23. Pessach, D., & Shmueli, E. (2022). *Algorithmic fairness*. *Communications of the ACM*, 65(7), 36–48.
24. Raji, I. D., Smart, A., White, R. N., Mitchell, M., & Gebru, T. (2020). *Closing the AI accountability gap: Defining an end-to-end framework for internal algorithmic auditing*. In *Proceedings of the 2020 Conference on Fairness, Accountability, and Transparency* (pp. 33–44).
25. Russell, S., Dewey, D., & Tegmark, M. (2015). *Research Priorities for Robust and Beneficial Artificial Intelligence*. *AI Magazine*, 36(4), 105–114.
26. Saurabh, P., Kumar, A., & Chauhan, S. (2022). *Responsible AI in digital transformation*. *International Journal of Information Management*, 62, 102450.
27. Schwartz, L., et al. (2020). *State of AI Ethics Report*. Montreal AI Ethics Institute.
28. Stahl, B. C., Timmermans, J., & Flick, C. (2017). *Ethics of emerging ICTs: Identifying the ethical issues*. *Technology in Society*, 46, 21–28.
29. Strubell, E., Ganesh, A., & McCallum, A. (2019). *Energy and Policy Considerations for Deep Learning in NLP*. In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics* (pp. 3645–3650).
30. Taddeo, M., & Floridi, L. (2018). *How AI can be a force for good*. *Science*, 361(6404), 751–752.
31. Véliz, C. (2020). *Privacy Is Power: Why and How You Should Take Back Control of Your Data*. One World Publications.
32. Yang, Q., Liu, Y., Chen, T., & Tong, Y. (2019). *Federated machine learning: Concept and applications*. *ACM Transactions on Intelligent Systems and Technology (TIST)*, 10(2), Article 12.