

International Journal of Advance Research in Computer Science and Management Studies

Research Article / Survey Paper / Case Study

Available online at: www.ijarcsms.com

A Monthly Double-Blind Peer Reviewed, Refereed, Open Access, International Journal - Included in
the International Serial Directories

Ethical AI versus Dark Patterns: Conceptualising a Trust- Based Framework for Marketing Automation

Rajesh Poonia¹

Faculty of Commerce and Management
SGT University,
Gurugram, Haryana, India.
rajeshpoonias@gmail.com

Jyoti Godara²

Faculty of Engineering and Technology
SGT University,
Gurugram, Haryana, India.
jyotipoonia6@gmail.com

DOI: <https://doi.org/10.61161/ijarcsms.v13i6.4>

Received: 17 May 2025; Received in revised form: 05 June 2025; Accepted: 15 June 2025; Available online: 18 June 2025

©2025 The Author(s). Published by IJARCSMS Publication. This is an open-access article under the CC BY license

(<https://creativecommons.org/licenses/by/4.0/>).

Abstract: *The emergence of algorithmic systems in digital marketing has boosted the prominence of automation in the Digital marketing sphere to the frontline in consumer engagement programs. Where the modern technologies of marketing are based on AI, and are intended to enhance personalization and make the user experience complete, recent research points towards the increase of the so-called dark patterns, or interface designs that covertly manipulate consumer choices towards corporate goals. Among them are coercive concurring creeks, blackened opt-out channels, and cheating calls to action, which carry along psychological and emotional predispositions (Brignull, 2010). The current discourse builds a theoretical model that allows differentiating between ethical and manipulative automated strategies in marketing environments. With a reminder of previous literature on trust and transparency, as well as trust towards algorithms and machine learning (Pasquale, 2015; Raji et al., 2020), the framework proposes the concept of Trust-Based Marketing Automation (TBMA) that aims to find a balance between optimization in favor of efficiency and consumer-centric prioritization of fairness. It is argued that dark patterns can create short-term gains in terms of conversions, but that they will co-occur with long-term losses to trust in customers and brand equity (Luguri & Strahilevitz, 2021). The study has presented four interconnected constructs of perceived fairness, algorithmic intent, user autonomy and interface transparency through an interdisciplinary analyst skill that defines ethical boundaries of marketing automation. The conceptual model, in turn, can lead to the ongoing scholarly discourse of the concept of responsible AI in consumption contexts and provide business management with recommendations on how to design trustful digital experiences.*

Keywords: *Ethical AI, dark patterns, marketing automation, consumer trust, algorithmic transparency, persuasive design, digital ethics.*

I. INTRODUCTION

The spread of artificial intelligence (AI) in the marketing field has transformed the relationship between firms and consumers because it makes the customization of the content automated, the delivery of the content more precise, and the ability

to target consumers in real-time and at mass scale. Modern marketing automation systems--part and parcel of customer relationship management (CRM), online advertisement, and customer personalization in online shops--are becoming increasingly dependent on AI to facilitate more efficiency and relevance in communication operations (Kim & Dennis, 2019; Choi, Kim, & Shin, 2023). However, such technological developments also create a substantial amount of ethical questions. They can utilize cognitive biases, inhibit the freedom of consumers, and weaken the boundary between persuasion and manipulation as the very operations that result in personalizing their content (Gray et al., 2018; Luguri & Strahilevitz, 2021; Budshra et al., 2024; Mohan, 2021).

At the center of this objection, there is the increasing trend of using dark patterns--a design practice that misleads or nudges users into making decisions, which they would otherwise not be keen on (Brignull, 2010; Mathur et al., 2019). These tactics, adopted by the modern robo-marketer as the de facto strategy of personalization (or conversion maximization), involve hidden opt-outs, deceptive indicators of urgency and other schemes that are detrimental to active consent. To a great extent, based on information and strengthened by algorithmic testing, their transparency makes them hard to trace by consumers and regulators alike (Narayanan et al., 2020; Kahng, Kim, & Sundararajan, 2022). In the growing awareness of marketing automation systems, there will be a burning need to distinguish between ethical AI use and unscrupulous design processes that undermine consumer confidence and disrupt the emerging digital rights standards (Kumar et al., 2024; Malhan & Kumar, 2021; Pankaj et al., 2023).

This paper proposes a conceptual model referred to as trust-based marketing automation (TBMA) that can be used to describe the extent of safe use of technology and technological manipulation to enforce compliance. TBMA emphasizes the significance of transparency, informed consent, and the agency of users as the key features of appropriate ethical practice, thus allowing the future research and the regulatory community to track the dangers of the AI-driven marketing communication and prevent them.

The ethical implications of digital persuasion in marketing and human-computer interactions (HCI) have been discussed with extant literature, although the discussion in marketing and HCI is fragmented, with the discussion organized by a topic, e.g., privacy (Martin et al., 2017), interface design ethics (Gray et al., 2018), and algorithmic fairness (Kumar, 2021; Kumar et al., 2023) (Binns et al., 2018). The existing models of digital marketing largely emphasize functional utility (efficiency, or personalization or lead conversion), and fail to pay attention to how those systems affect the levels of perceived fairness, autonomy and -trust among consumers (Shin, 2021; Awad & Krishnan, 2022). In the same respect, AI ethics frameworks promote other values, including transparency and explainability, but rarely do they apply those concepts to guidelines that can guide ethical design in marketing settings (Floridi et al., 2018; Raji et al., 2020).

Even though the discussion of dark patterns is currently gaining traction (Luguri & Strahilevitz, 2021; Mathur et al., 2019), there is little conceptual integration with broader psychological concepts, such as consumer trust, autonomy, and fairness. With more companies applying automated decision making in their digital interface space, there is an urgent infrastructure that needs to be developed that extends beyond the legal violations, and instead creates models that can be built that can help determine when marketing automation becomes ethical and when it is manipulative and coerces (Kumar et al., 2025; Yadav & Sahoo, 2025; Nguyen et al., 2025; Yadav et al., 2024). This gap will be addressed in the current paper that will propose the Trust-Based Marketing Automation (TBMA) framework that combines concepts of marketing ethics, AI design, and digital psychology to outline when an automated solution can be seen as ethical, transparent, and trustworthy. The framework does not only organize the key factors that influence the consumer trust into automated interfaces but also delivers a set of testable hypotheses that can be used to study the topic in the future empirically.

II. CONCEPTUAL BACKGROUND

The current paper proposes a multi-faceted theoretically-informed model that explains how AI marketing automation platforms can be seen as ethical, equitable and trust-building under individual circumstances. To achieve this, the article

suggests the Trust-Based Marketing Automation (TBMA) framework that defines and connects six fundamental constructs: algorithmic intent, interface transparency, user autonomy, perceived fairness, customer trust and ethical marketing automation outcomes. These constructs are the building blocks to a more human approach to AI marketing, shifting the attention away from manipulative controls of behaviors and decisions towards persuasion as focusing on the agency and dignity of the users.

It implies that the TBMA framework incorporates the knowledge of several fields, which should provide a solid conceptual framework. Using the literature on marketing ethics and relationship marketing, it will recognise that trust is a key ingredient which cannot be ignored in long term involvement of consumer and brand relationships (Morgan and Hunt, 1994; Mayer et al., 1995)(Kumar et al., 2024; Mohan, 2021; Sushma et al., 2025). It borrows concepts and findings in HCI and the persuasive systems design study, such as the fact that interface cues and patterns can and are used to shape user behaviours, as well as indicate dangers of dark patterns taking away autonomy and transparency (Bosch et al., 2016; Utz et al., 2022). Last, it relies on AI ethics and algorithmic accountability, replacing the reliance on a system design based on explainability, fairness and responsible use of data as crucial components of ethical system design characteristics (Floridi et al., 2018; Mittelstadt et al., 2016).

The suggested approach to triangulating the TBMA framework enables viewing TBMA as the disciplinary bridge that responds to the need to generate interdisciplinary models that guide theoretical applications as well as practical design choices in AI-enabled marketing. With references in the study to existing psychological and ethical framework, as well as consideration to newer regulatory and design issues, the study poses TBMA as a prospective solution to the crisis of trust that the modern marketing automation is facing.

2.1 Algorithmic Intent

Algorithmic intent A concept of research on algorithmic governance refers to the perceived motive or aim in automated systems, regardless of whether those systems were developed with consumer welfare in mind as their key aim, or whether they are endangered by commercial pressures in the guise of personalization. Since the nature of AI must reflect similar aspirations and values of those who create it, algorithmic intent can take the shape of ethical personalisation or an insidious form of control. It is hard, though, to decide which motives lie behind the algorithmic nudges: whether it is moving closer to the interests of the user or the exploitation of the behavioural weaknesses of people, as the choice of complex black-box decision systems can hinder the identification of individual user agency.

Empirical sources are replete with consideration that commercial AI is unduly optimized in the direction of engagement, internet keeping and conversion numbers, instead of admirer welfare (Zuboff, 2019). Such orientation drives algorithmic reasoning in the direction of exploitation in environments that involve using micro-targeting and behavioural data in tandem (Eubanks, 2018). In AI-enhanced personalisation in marketing, the boundaries can be crossed into coercive manipulation, as these measures prompt people to make snap judgments and fall into the trap of not showing how the nudge works (Sunstein, 2015; Narayanan et al., 2020). Subsequently, the motive on which automation is intended, is a determinant factor that motivates researchers to ground it on its ethics. As soon as the algorithm is planned to promote the interests of the firm on the pretext of supporting the consumer, it is venturing into the space of a dark pattern (Mathur et al., 2019). Comparatively, the approach to marketing automation may be considered morally accountable in the case of fully visible consumers reproduction of algorithm protocol (when it is in line with consumer objectives) and when the companies are open regarding trade-offs and granting users with real preferences (Mittelstadt et al., 2016; Raji et al., 2020).

The Elaboration Likelihood Model (Petty & Cacioppo, 1986) is a commonly used theoretical tool. Through this, it asserts that thinking in the context where the consumer evaluates the attempts to influence them as being manipulative or unreal makes them to mobilize counter-arguments and develop cognitive resistance. Luguri & Strahilevitz (2021) state that such algorithms that conceal or are deceptive erode persuasion and the long-term trust. Therefore, modern ethical systems of AI

advocate the idea of introducing intentional transparency into the process, where the purpose of automation whatever it is sales, satisfaction, or social good must be stated clearly (Floridi et al., 2018).

At the same time, the works on algorithm auditing show that the intent to cause harm or exploit can be reverse-engineered. Examples of such tools include SHAP and LIME tooling, which can demonstrate the features that were considered most important in the prediction process, thus enabling investigations of whether ethical personalization (e.g., surfacing relevant items) or exploitative manipulation (e.g. by introducing a sense of scarcity or so-called dark UI) has been programmed into the AI design.

Empirical literature further reveals that the idea of the perceived intent of algorithms has a significant impact on user trust and their attitude of compliance. When consumers feel the use of an AI system is serving their interests such as suggesting healthier foods as opposed to more popularized ones then they feel better in terms of their satisfaction, their perception of fairness and acceptance of automation in general (Lee & See, 2004; De Vries et al., 2020).

Considering that, this paper suggests that Algorithmic Intent can become a primary construct in the new field of Trust-Based Marketing Automation (TBMA). It draws the line between the influencing intentionally and covert manipulation claiming that ethical marketing automation should include the algorithms aimed at empowering people instead of tricking them.

2.2 Interface Transparency

Interface transparency deals with how well a marketing interface discloses, with clarity, openness, and access, its functionality, data practices and decisions to the user. It corresponds to the extent to which the users can make out what the system is up to, and why. In the situation of AI-enabled marketing, interface transparency has become more acknowledged as a premise of ethical interaction design. The shrouding of vital information by, e.g., auto-selected options, concealed charges and opt-in default mechanisms fall into the broad category of dark patterns, as developed by Brignull (2010). These user exploitative designs prevent informed decision-making and suppress user agency resulting in a low level of trust (Mathur et al., 2019).

The implication of historical research on transparency is that it forms the basis of trust within the digital system. The HCI literature states that transparency of interfaces may take three forms: operational transparency (how the system works), data transparency (what is collected), and intent transparency (why some particular choices are suggested) (Ananny & Crawford, 2018). The users are unable to determine fairness or integrity of the system when any of such layers is lacking (Burrell, 2016).

Empirical research indicates that due to their increasing perceived fairness and reducing suspicion, transparent interfaces tend to be perceived as more appropriate in the context of algorithmic decision-making, where decision rules are usually construed as opaque (Wachter et al., 2017). The users are more engaged and satisfied when presented with cues of explanation (e.g., You are seeing this ad because...). (Diakopoulos, 2015). On the other hand, perceived deception is linked to lack of transparency, especially in cases where users find that defaults were hidden or came to awareness of having accidentally agreed to do something on a hindsight (Luguri & Strahilevitz, 2021).

As a normative pillar, the General Data Protection Regulation (GDPR) implies transparency so that citizens are aware of the personal data collected and how them are processed (Voigt & Von dem Bussche, 2017). Nonetheless, empirical evidence suggests that companies commonly obscure these disclosures in lengthy privacy statements or use friction on the interface disincentivizing opt-outs (Nouwens et al., 2020). These are practices that reduce other than creating transparency hence we can consider them as ethics violations.

The researchers promote the concept of a designing in transparency paradigm; consequently, the interactive system directly incorporates real-time feedback, modular user-specified consent processes, and explanations that should empower the end user (Shneiderman, 2020). This view finds parallel in a recently emerging research area known as explainable AI (XAI),

which entails the development of systems that provide understandable explanations of the results generated by algorithms such that the end-users have a grasp of relevant decisions (Doshi-Velez & Kim, 2017).

In marketing automation processes--individualized email prospecting, introductory subscriptions and advertising placements on social media such as Facebook, Twitter or Instagram--interface transparency is of special interest. They are all criticized when it comes to such practice as prolonging trial periods and transferring them into subscriptions through confusing texts, when it comes to interfaces that do not use obscure texts and that highlight the possibility of canceling or rejecting instead, customers are more likely to trust them (Gray et al., 2018).

An increased empirical evidence base indicates that clear interface raises the perceived usability, lowers cognitive loads, and strengthens the brand perception (Eslami et al., 2018). Additionally, transparency is a catalyst mechanism: the more these users feel the fairness, the stronger their autonomy and the associated trust (Harrison et al., 2020). In line with this, in the framework of the Transactional Brand Management and Advantage (TBMA), Interface Transparency has been established as an organizational milieu, which promotes both the technical construction and the psychological impact of relations between consumers and brands.

2.3 User Autonomy

The meaning of user autonomy in marketing automation: A measure of the extent to which consumers feel they have real control over own preferences without being coerced or lied to, or being manipulated through algorithm twitching, or any manipulation. This construct plays a central role in the ethico-inquiry in an AI-guided marketing environment, referencing the ability of a consumer to effectively use self-governing judgment without hindrance by a presumably veiled manipulative approach. In dark patterns, theoretical approaches identify that consumer choice in interface design is violated systematically: interface design does not afford the consumer a sufficient range of consumer options, or even causes consumers to act in ways they do not desire (Luguri and Strahilevitz, 2021).

According to empirical research studies, dark patterns manipulate these cognitive biases and behavioral inertia (including such as default bias, scarcity, urgency and confirmation bias) to overcome rational consumer preferences (Gray et al., 2018; Mathur et al., 2019). Such manipulations are not only forceful sales approaches but vulnerable practices that break down consumer sovereignty and, in most cases, breach the ethical value of informed consent (Kahng et al., 2022). Psychologically, autonomy is one of the key components of Self-Determination Theory (Deci and Ryan, 2000), and argues that the behaviors people feel most motivated and at their satisfaction level are self-endorsed. User control interfaces that coerce or make users make a decision (i.e., force them to share their data, subjecting them to the dark side of the opt-out process) are felt as undermining autonomy and generate opposition and discontent (Utz et al., 2022). There is also the increasing constraint on user autonomy by the use of digital-marketing methods that can be described as roach-motel designs, which are the structures where the user can easily subscribe to but face a lot of challenge to cancel the subscription (Mathur et al., 2021). Such imbalances, which are common in the software-as-a-service scenario, do economically benefit corporations at the expense of consumer welfare. Recent research on consumer advocacy groups and algorithmic regulation agencies have concluded that the patterns of coercion should be defined as violations of the new regulatory frameworks, e.g. the California Privacy Rights Act (CPRA) and the EU Digital Services Act (Eidelman et al., 2023).

To this end, the current champers of ethical marketing espouses autonomy-sensitive approaches placing emphasis on just-in-time permission, adequate explication, irrevocable choice, reverse requesting systems (Kelley et al., 2022). Tracking preferences or any other process of personalization logic given by the users can employ interfaces that enable them to feel empowered and is associated with satisfaction and long-term brand engagement (Yao et al., 2021). Convergent evidence in persuasive-computing areas suggests that consumer trust is not the only positive effect that autonomy-encouraging designs have on compliance with ethical standards, especially in delicate areas like healthcare and financial advertising (Koencke et al.,

2023). This is due to treating the user with dignity and matching interventions with values rather than acting in a sneaky manner of influencing the choice. In cases where autonomy is seen to be secure, users are likely to describe brands as transparent and equitable and socially responsible (Huang & Bashir, 2021).

The recent trend in AI-generated content, customizable persuasive agents, e.g. conversational bots, recommender systems, brought new questions to the field of autonomous agency. Consumers may perceive automated nudges as genuine advice in the instance of machine models rehashing human persuasive strategies (e.g., pertinence warnings in conversational bots), eroding their abilities to question intentionality (Lee et al., 2023). The existence of such deception-by-design highlights the need to have theoretical paradigms that position autonomy as the center piece of ethical AI marketing. In this analytical framework, User Autonomy thus goes beyond the issue of interface choice, being an ethical duty of cognitive independence in face of algorithmic persuasion.

2.4 Perceived Fairness

Fairness is a judgement by a consumer of whether an AI-assisted marketing does the right things, correctly, and without discrimination through all processes of decision making, interface design, and data-management practices. Algorithmic fairness has become one of the central issues in the field of responsible AI, especially when it comes to the scenarios of trial-and-error deployment of AI technologies in the areas of finance, healthcare, or advertising (Wachter et al., 2021; Lee & Singh, 2023). In the marketing environment, algorithmic fairness perceptions enter the picture and shape the customer assessment of the very integrity of the brand-customer relationship, in addition to particular offers or personalisation strategies (Kahng et al., 2022).

Consumer research has long separated three dimensions of fairness: the fairness of outcomes, the fairness of procedures and the fairness in communication and explanation (Tax et al., 1998). These differences are being transferred to AI environments in which customers are also seeking fair results in addition to algorithmic fairness themselves algorithms that do not engage in exploitation or biased personalisation (Binns, 2018). The persisting discussions of algorithmic discrimination, especially those relating to pricing functions, content delivery, and credit scores, confirm the high relevance of explainability and transparency as the factors promoting the sense of fairness (Binns et al., 2018; Sandvig et al., 2014). Such a tension is specifically captured by dynamic pricing algorithms that charge different users a different price based on opaque profiling; such systems can legally exist, but are perceived as unfair to consumers, resulting in negative brand perceptions, regulatory intervention, and subsequent backlash.

Empirical studies however indicate that perceptions of fairness acts as a connecting factor between personalization and customer satisfaction (Lee & Shin, 2022). Although personalization is listed as a positive activity in ethically designed systems, the opaque case of segmentation or exploitive targeting weakens trust and satisfaction because it reduces the cues associated with fairness (Kahng et al., 2022; Kim et al., 2021). Social comparison theory also contributes to the evaluation of fairness by consumers: people often consider fairness based on what results they receive in relation to what others can receive in similar conditions (Festinger, 1954). In modern systems of digital sales, where discount offers, targeted with specific advertisements, and different subscription plans make such disparities extremely visible, the sense of inequity becomes one of the significant variables predicting turnover and dissatisfaction (Huang & Bashir, 2021). On operations, the maximization of fairness will be most efficient in case firms provide opt-out controls, clear justifications of recommendations, and a system transparency that allows users to interpret factors underlying reasoning (Yao et al., 2021). These measures call upon interactional fairness and procedural fairness as they also state that the user has the right to know and decide. It has been observed empirically that the use of explanation-based transparency significantly increases the fairness considerations, irrespective of the actual outcomes (Binns et al., 2018; Harrison et al., 2020). One of the recent literature streams has laid a construct of perceived data fairness as the consumer perceived that data collection and use is undertaken in a respectful proportional manner (Kahng et al., 2022; Eslami et al., 2018). It is revealed that these worries about the amount of data collected and the danger of misuse are producing anxiety rooted in fairness factors in technical safe systems. As the TBMA table shows, Perceived Fairness plays an only mediating role

between the technical features of interface (transparency, autonomy, etc.) and psychological outcomes (trust, satisfaction, and so forth). Without fairness, even the best systems may be taken as being manipulative or unfair, thus, strengthening the suspicion of AI-powered marketing automation.

2.5 Customer Trust

Customer trust refers to a psychological reliance, wherein consumers avail of the system of marketing automation of an organisation at their own freewill, as they consider them to be capable, trustworthy, and in line with user interests in situations where AI comes into the picture. Trust is a component in modern marketing which forms a core building block of long-term customer relationships specifically in cases where a technology presents a mediation agent in the process. Authors argue that as AI increasingly is used in personalisation, targeting and automation, the foundation of trust has been changed to a system-based trust fueling the necessity of the consumers to trust algorithms they do not always comprehend (Lankton et al., 2015; Belanche et al., 2020).

AI-driven marketing automation trust is a multidimensional construct based on competence (does the system act accurately), integrity (does the system act in a fair and honest manner), and benevolence trust (acts in the interest of the customer) (Mayer et al., 1995; McKnight et al., 2011). A violation of one of these attributes, e.g. abuse of data, inequitable targeting, and manipulative nudges, is to impair the trust of users severely, potentially leading to reputational loss (Luguri & Strahilevitz, 2021; Wirtz et al., 2022).

Most modern sources emphasize the frailty of algorithmic-based trust, especially in cases when the explainability is low, or people feel no control over the outcomes (Shin, 2021). In contrast with human actors, AI agents do not manifest emotional signals or specific reputations and, therefore, the degree of trust during human-computer interactions should be achieved by means of interface transparency, ethical personalisation, and procedural fairness (Eslami et al., 2018; Awad & Krishnan, 2022).

Glikson and Woolley (2020) prove that users demonstrate the willingness to update the level of trust they have in the AI-generated recommendations toward the higher end of the scale in case the system proves itself fair and competent in its suggestions on multiple occasions. This finding corroborates other facts that AI trust is dynamic and varying depending on the system design and ethically good performance (e.g. Koenecke, Stegen, et al., 2023).

Trust in automation in the literature has been related to brand loyalty, reduced customer churn and increased likelihood to share personal data (e.g. Lankton, Nguyen, et al., 2021). On the contrary, when the limits of AI-induced interactions ethical scope, or when the dark patterns are revealed (e.g. automatic subscription, false urgency, hidden costs), users feel betrayed; they react with distrust and, despite the overall positive brand image, it becomes less reliable (Gray et al., 2018; Mathur et al., 2019).

The theory of trust transfer gives additional knowledge, as it provides examples of how the trust in a piece of technology (its interface) may be transferred onto more general assessments of brand. In cases where automated touchpoints, such as email communications or chatbots, recommendation engines, etc., are found to be trustworthy, ethically good, and responsive, customers tend to generalize a set of these positive features over to the brand as a whole (Wirtz et al., 2022; Choi et al., 2023). Conversely, adopting unethical design approach is capable of destroying the perceived trust in a system and brand.

Trust is an emerging regulatory imperative within the scholarship of AI ethics, with a number of organizations and initiatives like the OECD, the IEEE, and the EU AI Act promoting the concept of a trustworthy AI as a fundamental goal of design (Floridi et al., 2018; European Commission, 2021). The message of this call is clear to marketers, and that is that automation practices need to converge with transparency, accountability, and consumer agency to enhance trust as a strategic resource and a competitive edge. Within the TBMA framework, Customer Trust acts as an outcome of upstream variables like algorithmic intent, transparency, autonomy, and fairness. When these inputs are ethically aligned, trust emerges not just as an emotion but as a rational confidence in system integrity—making it indispensable to sustainable marketing automation.

2.5 Ethical Marketing Automation Outcomes

Marketing systems with its key values as transparency, fair treatment, agency, trust, and ethical AI design result in consumer responses which take time but consist in loyalty, satisfaction, advocacy, and decreased churn. With modern algorithmic marketing, it is possible to have evaluative processes that are based on an assessment of short-term metrics like click through rates, open rates and conversion. However, the current studies prove that ethical-based offerings provide more long-term, trust-based results that go beyond transactional KPIs (Wirtz et al., 2022; Kim & Dennis, 2019). Relationship quality as opposed to quantitative engagement marks the ethical aspect of marketing automation.

The empirical evidence reveals that responses of consumers to marketing automation are not merely determined by the relevancy of content but rather by the perceived need of ethical alignment, including data fairness, user agency, or adherence to personal preferences (Martin et al., 2017; Choi et al., 2023). Ethical AI is an indicator of the appreciation of the dignity of consumers, which strengthens the customer-brand relationship (Shin, 2021). On the other hand, there is growing sensitivity among the consumers as to the dark UX strategies and the fight against the manipulative automation is becoming stronger and stronger (Luguri & Strahilevitz, 2021). The studies found that perceived manipulative practices, such as forceful continuity, bait-and-switch system, and coerced consent, are concurrent to low satisfaction, trust compromise, and political word-of-mouth (Mathur et al., 2019; Kahng et al., 2022). In comparison, ethical automation platforms foster loyalty and propagating behavior, particularly when compatible with consumer values, like privacy, honesty, and personalization transparency (Awad & Krishnan, 2022).

In the most recent study published by Kim, Jung, and Sundar (2022), it empirically demonstrated that, in some cases, in which the predictive accuracy is somewhat declining, the more transparent the intentions of AI-based personalizing systems are articulated and the easier the opt-out mechanism is provided to users, the more favorable their attitude becomes. This observation supports the idea that ethical friction, in the form of the existence of consent frameworks and perceptions of user control, can promote satisfaction under the conditions in which it is also perceived as ethical in regards to voice or autonomy. Moreover, Sinha and Singh (2022) reveal that, indeed, customer-friendly AI design, with clear explanations, real-time click-level changes, and the possibility of registering feedback, has a strong impact on the qualities of brands, purchase intents, and trust in the virtual world. When privacy by design is a legal requirement, the companies that incorporate ethics in the automation decisions gain competitive benefit, such as in Europe and in California. Similarly, Martin et al. (2021) and Floridi et al. (2018) claim that notification of fairness, transparency, and empowerment of consumers increases ethical consequences upon adopting them in messaging and interface design. Automation becomes an organization-building process when the organizations have incorporated such practices as opposed to its use as a persuasion tool.

Collectively, all these studies demonstrate that in the TBMA framework, Ethical Marketing Automation Outcomes are the functional test of ethical design and AI responsibility as applied. They illustrate the downstream effects of the algorithmic intention, clarity, independence, equitability, and confidence to be desirable consumer actions; that is, high levels of customer loyalty, augmented perceived brand genuineness, less customer complaints, and amplified customer satisfaction and brand loyalty. Instead, ethical marketers no longer define marketing effectiveness based on immediate results (e.g., clicks), but switch to an aim of creating sustainable value and treating each other with respect, which is ultimately the future potential of responsible AI in marketing.

III. HYPOTHESIS AND THEORETICAL MODEL

H1: Algorithmic Intent has a significant positive effect on Perceived Fairness.

H2: Interface Transparency has a significant positive effect on Perceived Fairness.

H3: Interface Transparency has a significant positive effect on User Autonomy.

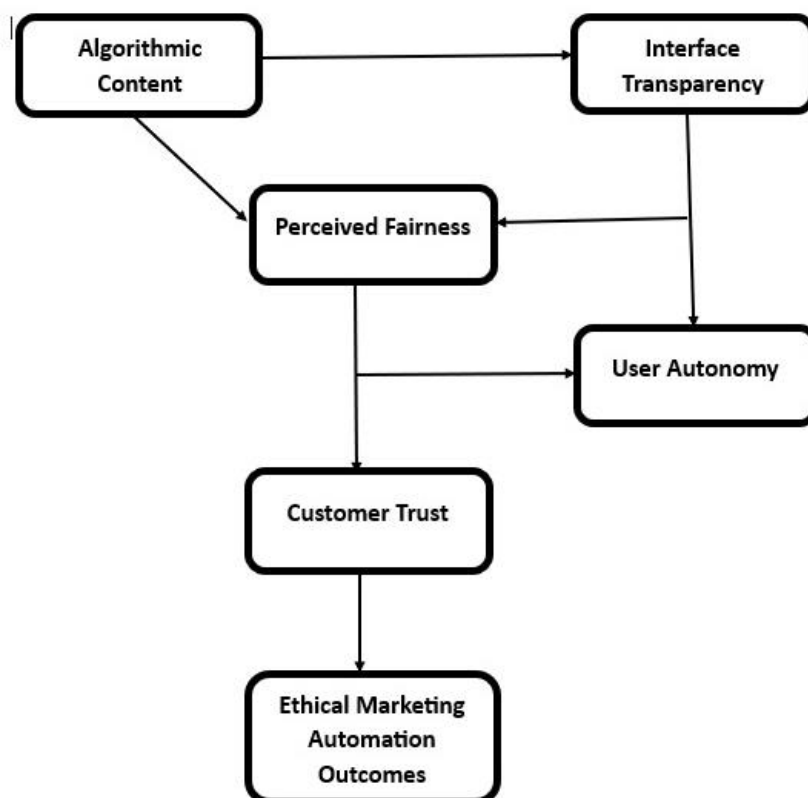
H4: Perceived Fairness positively influences Customer Trust.

H5: User Autonomy positively influences Customer Trust.

H6: Customer Trust positively affects Ethical Marketing Automation Outcomes.

H7: User Autonomy positively affects Ethical Marketing Automation Outcomes.

H8: Perceived Fairness positively affects Ethical Marketing Automation Outcomes.



Theoretical Model: Trust-based Marketing Automation

This study contributes a theoretical and conceptual framework, which builds on an interdisciplinary synthesis of the knowledge on marketing ethics, human-computer interaction (HCI), and algorithmic design. The most important is to develop Trust-Based Marketing Automation (TBMA) framework, a conceptual tool that will explain the interpretive perception and response of consumers to the AI-enabled systems that either practice ethical or manipulative processes.

The TBMA system is built on six key constructs: the intent of algorithm, transparency on interface, autonomy of the user, the perceived fairness, the customer trust, and the consequences connected with an ethical automation of the marketers. These components have no relation to empirical observation but are based on the thorough overview of the scientific literature on algorithmic persuasion, personalization ethics, dark patterns, and consumer trust. All the constructs are chosen due to their frequent appearance in the modern discussions of the topics. The model is based on the assumption that the marketing automation systems designed with ethically-oriented principles, when seen as transparent, fair, and autonomy-enhancing, can give way to more envisionable and trust-based consumer outcomes.

The TBMA theory of the framework has intersections in many areas. Trust, fairness, and relationship quality are constructs that can be found in the marketing and behavior science classes but they were broadly discussed in the literature on consumer-brand engagement and CRM (Mayer et al., 1995; Lankton et al., 2015). To accompany these insights, HCI and UX design contribute such concepts as interface manipulation, transparency layers, and cognitive friction (Gray et al., 2018; Binns et al., 2018). Lastly, the framework relies on some arising issues in the field of AI ethics, and digital regulation such as explainability,

accountability of algorithms, consumer protection in automated decision-making scenarios (Floridi et al., 2018; Raji et al., 2020).

The current work is not aimed to measure constructs in a scientific way or prove the stated hypotheses with the help of statistics. Instead, its major contribution should be seen in terms of conceptual integration, systematically cross-aligning disparate awareness in different fields to provide an internally coherent template by which future empirical studies can challenge the trust-oriented design of automation. The hypotheses presented in this paper are deliberate in a manner geared towards guiding future investigations and may easily be tested through various research strategies; they include:

- Empirical verification with the use of surveys that usually use structural equation modeling (SEM) or partial least squares (PLS) to evaluate the relations between constructs;
- Ethical design of the interface contrasting the reaction of users to ethical and non-ethical interfaces involving dark patterns;

Qualitative tests of consumer stories, beliefs in algorithmic fairness or experiences of trust disintegration caused by manipulative design. With the help of providing a theoretical scenario guided by various approaches, the suggested model will be able to guide the research process as well as the managerial behaviour. In this way, it contributes to the discussion of ethical AI in marketing and opens the gate to more people-centred and regulated future of automated persuasion.

IV. IMPLICATIONS AND FEATURE RESEARCH DIRECTIONS

The Trust-Based Marketing Automation (TBMA) framework presented in this paper offers both theoretical contributions and actionable implications for scholars, marketers, interface designers, and policy-makers navigating the ethical challenges of AI-driven consumer engagement.

4.1 Theoretical Implications

The given research makes a contribution to a range of theoretical perspectives. To begin with, it brings the constructs of marketing, behavioral science, and computer science which have traditionally been researched in separate domains to unite them into a set of constructs that explain how ethical aspects influence consumer impressions about automation. Instead of approaching the classification of AI as either ethical or unethical, the model provides the necessary multilayered framework within which the concept of trust can serve as the key mediator between the technical aspects of interface design and the eventual consumer-long-term consequences.

Second, the framework enhances current theories of trust formation and fairness amongst consumers since concerning the intention of algorithms and interface clarity, which have been largely ignored in traditional marketing research yet are key to the digital experience in AI systems. As a result, the TBMA model can facilitate the transition between marketing ethics and AI governance and, thus, comply with the emerging demands to develop conceptual frameworks capable of combining the psychological understanding of ethical issues with the technical facts of life.

Third, the model outlines testable hypotheses that can be verified in different studies by future scholars who choose to conduct empirical research using a quantitative or a mixed-methods study. It sets out a research agenda, which is cumulative, in regard to the ethical utilization of AI applications in marketing.

V. MANAGERIAL AND POLICY IMPLICATIONS

The current paper provides a strict theory of marketing automations in terms of measuring them and designing them without jeopardizing, yet increasing consumer trust. In a world where there is a plethora of personalization and behavioral targeting, the TBMA model triggers the marketers to break out of click-based strategies to more responsible and user-based strategies. The research therefore has four managerial and policy implications, namely:

- Design for transparency: Brands have to give clear, understandable notifications on how the AI systems work and what kind of data they use and on what basis some certain decision is made.
- Respect user autonomy: Whether by marketing interface or opt-in/opt-out default, marketing interfaces should be allowed to be coercive, as in default opt-in offers should at least not be imposed and cancellations should not be made otherwise difficult. Data-handling and consent decisions should be free and significant.
- Embed fairness safeguards: Algorithms should do nothing to systematically discriminate at particular consumer segments or produce discriminatory outcomes.
- Monitor algorithmic intent: System design should be focused on the well-being of consumers other than on the performance indices. This congruency must be reflected in front-end view as well as back-end logic.

Regulators and policymakers can use the TBMA framework as their basis in terms of setting ethical standards in marketing automation, as it aligns with the legal regulatory acts, policies, such as the EU Artificial Intelligence Act, GDPR, and OECD Principles on Trustworthy AI, simplifying their implementation on the interface level.

VI. FUTURE RESEARCH DIRECTIONS

The present paper identifies a first conceptual territory, but it can be narrowed and extended in a variety of ways under further inquiry:

- Empirical validation: The strength and the direction of the hypotheses can be validated with the help of quantitative methodologies; it is possible to use surveys and structural equation modeling (SEM).
- Experimental studies: Lab or field experimental studies could be used to question the behavior and the processes of trust formation of users of various types of interfaces, such as transparent or dark-pattern interfaces.
- Contextual exploration: There is a reason to explore how the framework can be used in different sectors like health-tech, fintech, ed-tech and retail since the stakes and regulation environments will vary and tend to be heterogeneous.
- Cultural views: It is necessary to study the variation of perceptions of digital trust and fairness across different cultures and diverse demographical conditions.
- Longitudinal studies: Longitudinal studies should be informative in measuring the effect of good ethical course of automation practices on customer retention, loyalty and lifetime value.
- Intervention design: Their employment of the framework to develop and test its ethical design principles or toolkits that could help guide organizations into a responsible transition to marketing automation is extremely applicable.

Current study also provides a research question that is pressing and urgent: how can marketing automation systems be developed in a way to be persuasive rather than manipulative, personalization rather than invasion of privacy, and automation that would not harm consumer confidence in it?

VII. CONCLUSION

The study at hand contributes to the recent discussions of the concepts of responsible AI in marketing by elaborating the theoretically sound and comprehensive Trust-Based Marketing Automation (TBMA) framework that helps us understand under which circumstances the marketing automation may erode the consumer trust or strengthen it. Based on the available literature in marketing ethics, human-computer interaction (HCI), behavioral science, and AI governance, the model specifies six interrelated constructs of algorithmic intent, interface transparency, user autonomy, perceived fairness, customer trust, and ethical marketing automation outcome. Both technical and psychological aspects of the engagement between a user and automated marketing system are covered because both were reflected in the amplitudes of constructs against empirical studies in

marketing practice and ethical theory. Notably, TBMA does not frame trust as an inevitable outcome of technological integration, but instead a diaphanous, negotiable variable which is the product of interface design, algorithmic logic, and organizational purpose. The framework emphasizes, therefore, the ethical requirements to create systems that facilitate end-user power, adherence to fairness principles, and generate transparency, in effect, the prerequisites of a sustainable long-term association between brands, and the consumer who belong to an ecosystem that is becoming automated. TBMA despite being conceptual offers a series of hypotheses and empirically testable concepts that make it easier to engage in empirical research in the future. Both its core target and the scope of its aim ultimately reach beyond the mere purpose of informing academic scholarship by providing a decision-support tool to marketers, designers, and regulators wanting to align the mutuality of innovation with that of integrity. Finally, it is the framework that contradicts the current trend of optimization at all costs. The present-day research demands the creation of intelligent and yet deliberate marketing systems that could convey the temptation without acting up fraudulently, automated without acting tricky and personalised without infringing upon the prestige of a user. With newer technologies and marketing abilities growing together with AI, these trust-based models are bound to become the fundament of building an efficient yet morally sound technology.

References

1. Agrawal, U., Mangla, A. and Sagar, M. (2016) 'Company-cause-customer: interaction architecture', Global Journal of Flexible Systems Management, Vol. 17, No. 3, pp.307–319.
2. Ananny, M., & Crawford, K. (2018). Seeing without knowing: Limitations of the transparency ideal and its application to algorithmic accountability. *New Media & Society*, 20(3), 973–989. <https://doi.org/10.1177/1461444816676645>
3. Awad, N. F., & Krishnan, M. S. (2022). Building trust in AI: The role of algorithm transparency and user control. *MIS Quarterly Executive*, 21(1), 31–49.
4. Binns, R. (2018). Fairness in machine learning: Lessons from political philosophy. *Proceedings of the 2018 Conference on Fairness, Accountability and Transparency (FAT)*, 149–159.
5. Belanche, D., Casaló, L. V., Flavián, C., & Schepers, J. (2020). Service robot implementation: A theoretical framework and research agenda. *The Service Industries Journal*, 40(3–4), 203–225. <https://doi.org/10.1080/02642069.2019.1672666>
6. Bösch, C., Erb, B., Kargl, F., Kopp, H., & Pfattheicher, S. (2016). Tales from the dark side: Privacy dark strategies and privacy dark patterns. *Proceedings on Privacy Enhancing Technologies*, 2016(4), 237–254. <https://doi.org/10.1515/popets-2016-0038>
7. Binns, R., Veale, M., Van Kleek, M., & Shadbolt, N. (2018). 'It's reducing a human being to a percentage': Perceptions of justice in algorithmic decisions. *CHI Conference on Human Factors in Computing Systems*, 1–14. <https://doi.org/10.1145/3173574.3173951>
8. Brignull, H. (2010). Dark patterns: User interfaces designed to trick people. Retrieved from <https://www.darkpatterns.org>
9. Burrell, J. (2016). How the machine 'thinks': Understanding opacity in machine learning algorithms. *Big Data & Society*, 3(1), 2053951715622512. <https://doi.org/10.1177/2053951715622512>
10. Budshra, S., Malhan, D., & Kumar, M. (2024). Analyzing Work Stress Indicators and Remedial Measures Among Faculty in Analyzing Work Stress Indicators and Remedial Measures Among Faculty in Higher Education Institutions. *Frontiers in Health Informatics*, 13(4), 1572–1588.
11. Choi, J., Kim, J., & Shin, D. (2023). Trust me, I'm a bot: The role of AI agents in fostering brand trust and consumer engagement. *Journal of Interactive Marketing*, 61, 1–17. <https://doi.org/10.1016/j.intmar.2023.01.001>
12. Deci, E. L., & Ryan, R. M. (2000). The "what" and "why" of goal pursuits: Human needs and the self-determination of behavior. *Psychological Inquiry*, 11(4), 227–268.
13. De Vries, H., Kühne, R., & Realo, A. (2020). Consumer responses to algorithmic persuasion: How algorithmic agents influence perceptions of persuasion intent and credibility. *Computers in Human Behavior*, 112, 106459. <https://doi.org/10.1016/j.chb.2020.106459>
14. Diakopoulos, N. (2015). Algorithmic accountability: Journalistic investigation of computational power structures. *Digital Journalism*, 3(3), 398–415. <https://doi.org/10.1080/21670811.2014.976411>
15. Doshi-Velez, F., & Kim, B. (2017). Towards a rigorous science of interpretable machine learning. *arXiv preprint arXiv:1702.08608*.
16. Eslami, M., Rickman, A., Vaccaro, K., Aleyasen, A., Vuong, A., Karahalios, K., ... & Sandvig, C. (2018). I always assumed that I wasn't really that close to [her]: Reasoning about invisible algorithms in news feeds. *Proceedings of the 2015 CHI Conference*, 153–162.
17. Eubanks, V. (2018). *Automating inequality: How high-tech tools profile, police, and punish the poor*. St. Martin's Press.
18. Eidelman, S., Kalra, A., & Chin, A. (2023). Dark patterns in the wild: Patterns of manipulation across popular online platforms. *Communications of the ACM*, 66(4), 62–70. <https://doi.org/10.1145/3572766>
19. European Commission. (2021). Proposal for a regulation laying down harmonized rules on artificial intelligence (AI Act). Retrieved from <https://artificialintelligenceact.eu/>
20. Festinger, L. (1954). A theory of social comparison processes. *Human Relations*, 7(2), 117–140.

21. Floridi, L., Cowls, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., ... & Vayena, E. (2018). AI4People—An ethical framework for a good AI society: Opportunities, risks, principles, and recommendations. *Minds and Machines*, 28(4), 689–707.
22. Floridi, L., Cowls, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., ... & Vayena, E. (2018). AI4People—An ethical framework for a good AI society. *Minds and Machines*, 28(4), 689–707. <https://doi.org/10.1007/s11023-018-9482-5>
23. Gray, C. M., Kou, Y., Battles, B., Hoggatt, J., & Toombs, A. L. (2018). The dark (patterns) side of UX design. *Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems*, 1–14. <https://doi.org/10.1145/3173574.3174108>
24. Harrison, G., Flood, D., & Duivenvoorde, B. V. (2020). Consumers and dark patterns in personalized advertising: Data protection and privacy implications. *International Data Privacy Law*, 10(1), 37–57.
25. Huang, G., & Bashir, M. (2021). Dark patterns, revisited: Descriptive taxonomy and analysis of deceptive design patterns in e-commerce platforms. *Journal of Consumer Policy*, 44, 671–696. <https://doi.org/10.1007/s10603-021-09489-z>
26. Kahng, A., Kim, H., & Sundararajan, A. (2022). Algorithmic nudges and consumer choice: The role of autonomy and explainability. *Journal of Consumer Psychology*, 32(1), 152–168. <https://doi.org/10.1002/jcpy.1231>
27. Kim, K., Lee, J., & Sundar, S. S. (2021). Fairness perceptions of algorithmic price discrimination: Understanding user concerns and policy implications. *Journal of Business Research*, 133, 276–286. <https://doi.org/10.1016/j.jbusres.2021.05.030>
28. Kim, Y., & Dennis, A. R. (2019). When AI recommends: How do consumers respond to recommender systems in retail? *Journal of Retailing*, 95(4), 476–487. <https://doi.org/10.1016/j.jretai.2019.10.002>
29. Kumar, A., Sunder, S., & Sharma, S. (2020). Personalization versus privacy: An empirical examination of the online consumer's dilemma. *Journal of Business Research*, 106, 233–243. <https://doi.org/10.1016/j.jbusres.2019.09.023>
30. Kumar, M. (2021). An Evaluation of Entrepreneurial Tendencies of State Universities ' Students of North India : an Empirical Study. *International Journal of Entrepreneurship* (Print ISSN: 1099 -9264; Online ISSN: 1939-4675), 25(4), 1–11.
31. Kumar, M., Dagar, M., Yadav, M., Choudhury, T., & Luu, T. M. N. (2025). Navigating the future: Challenges and opportunities in intelligent RPA implementation. In *Intelligent Robotic Process Automation: Development, Vulnerability and Applications* (pp. 393–420). IGI Global. <https://doi.org/10.4018/979-8-3693-4365-4.ch015>
32. Kumar, M., Jain, A., Mittal, A., Gera, R., Biswal, S. K., Yadav, M., Hung, T. H., & Priya Srivastava, A. (2024). Inclusion of Neural Networks in Higher Education: A Systematic Review and Bibliometric Analysis. *2024 4th International Conference on Innovative Practices in Technology and Management (ICIPTM)*, 1–6. <https://doi.org/10.1109/ICIPTM59628.2024.10563852>
33. Kumar, M., Yadav, R., Suresh, A. S., Singh, R., Yadav, M., Balodi, A., Vihari, N. S., & Srivastava, A. P. (2023). Mapping AI-Driven Marketing: Review of Field and Directions for Future Research. *2023 2nd International Conference on Futuristic Technologies (INCOFT)*, 1–5. <https://doi.org/10.1109/INCOFT60753.2023.10425169>
34. Kahng, A., Kim, H., & Sundararajan, A. (2022). Algorithmic nudges and consumer choice: The role of autonomy and explainability. *Journal of Consumer Psychology*, 32(1), 152–168. <https://doi.org/10.1002/jcpy.1231>
35. Kelley, P. G., Cesareo, M., & Cranor, L. F. (2022). Friction or freedom? Re-evaluating consumer autonomy in opt-out systems. *ACM Transactions on Human-Computer Interaction*, 29(4), Article 25. <https://doi.org/10.1145/3533026>
36. Koenecke, A., Varshney, K. R., & Chouldechova, A. (2023). Manipulation or motivation? Ethical trade-offs in persuasive algorithmic systems. *AI & Society*. <https://doi.org/10.1007/s00146-023-01615-2>
37. Lee, J. D., & See, K. A. (2004). Trust in automation: Designing for appropriate reliance. *Human Factors*, 46(1), 50–80. https://doi.org/10.1518/hfes.46.1.50_30392
38. Lee, M., & Shin, D. (2022). The role of perceived fairness in trust formation in AI-based service platforms. *Computers in Human Behavior*, 130, 107182. <https://doi.org/10.1016/j.chb.2022.107182>
39. Lee, Y., & Singh, R. (2023). Fairness expectations in algorithmic marketing: A cross-cultural exploration. *Journal of Interactive Marketing*, 61, 50–68. <https://doi.org/10.1016/j.intmar.2023.01.002>
40. Luguri, J., & Strahilevitz, L. (2021). Shining a light on dark patterns. *Journal of Legal Analysis*, 13, 43–109. <https://doi.org/10.1093/jla/laaa006>
41. Lankton, N. K., McKnight, D. H., & Tripp, J. (2015). Technology, humanness, and trust: Rethinking trust in technology. *Journal of the Association for Information Systems*, 16(10), 880–918.
42. Lankton, N. K., Wilson, E. V., & Mao, E. (2021). The impact of automation and agency on trust in artificial intelligence. *Information Systems Research*, 32(4), 1226–1245. <https://doi.org/10.1287/isre.2021.1017>
43. Lee, M. K., Kizilcec, R. F., & Lee, J. J. (2023). Persuasive AI agents and the erosion of user autonomy: Experimental evidence and design implications. *Computers in Human Behavior*, 143, 107714. <https://doi.org/10.1016/j.chb.2022.107714>
44. Luguri, J., & Strahilevitz, L. (2021). Shining a light on dark patterns. *Journal of Legal Analysis*, 13, 43–109. <https://doi.org/10.1093/jla/laaa006>
45. Martin, K. D., Borah, A., & Palmatier, R. W. (2017). Data privacy: Effects on customer and firm performance. *Journal of Marketing*, 81(1), 36–58. <https://doi.org/10.1509/jm.15.0497>
46. Martin, K. D., Murphy, P. E., & Singh, J. J. (2021). Fairness in consumer digital markets: A conceptual framework and research agenda. *Journal of Public Policy & Marketing*, 40(2), 159–173. <https://doi.org/10.1177/07439156211004990>
47. Mathur, A., Friedman, M., & Narayanan, A. (2021). Manipulative web design at scale: Patterns in consent dialogs and implications for regulation.

- ACM Human-Computer Interaction, 5(CSCW1), 1–31. <https://doi.org/10.1145/3449102>
48. Mayer, R. C., Davis, J. H., & Schoorman, F. D. (1995). An integrative model of organizational trust. *Academy of Management Review*, 20(3), 709–734.
 49. Mathur, A., Acar, G., Friedman, M. G., Lucherini, E., Mayer, J., Chetty, M., & Narayanan, A. (2019). Dark patterns at scale: Findings from a crawl of 11K shopping websites. *Proceedings of the ACM on Human-Computer Interaction*, 3(CSCW), 1–32. <https://doi.org/10.1145/3359183>
 50. McKnight, D. H., Carter, M., Thatcher, J. B., & Clay, P. F. (2011). Trust in a specific technology: An investigation of its components and measures. *ACM Transactions on Management Information Systems (TMIS)*, 2(2), 1–25.
 51. Mittelstadt, B. D., Allo, P., Taddeo, M., Wachter, S., & Floridi, L. (2016). The ethics of algorithms: Mapping the debate. *Big Data & Society*, 3(2), 2053951716679679.
 52. Morgan, R. M., & Hunt, S. D. (1994). The commitment-trust theory of relationship marketing. *Journal of Marketing*, 58(3), 20–38. <https://doi.org/10.1177/002224299405800302>
 53. Malhan, D., & Kumar, M. (2021). A Study on Measuring Impact of Social Media Marketing on Consumers' Buying Behavior Based on Demographic Variables. *Marketing 5.0: Putting Up Blocks Together* Publisher: National Press Associates, New, May, 125–130.
 54. Mohan, Nisha, D. D. M. D. P. (2021). Influence of perceived customer satisfaction on online shopping experience during the COVID-19 pandemic. *The Journal of Contemporary Issues in Business and Government*, 27, 1002–1012. <https://api.semanticscholar.org/CorpusID:239644725>
 55. Mohan Kumar, Yadav, M., & Sahoo, A. (2025). The Importance of Collaboration Skills in Workforce Preparation From Textbooks to Teamwork. In *Revitalizing Student Skills for Workforce Preparation* (pp. 333–352). <https://doi.org/10.4018/979-8-3693-3856-8.ch011>
 56. Nguyen, T. B., Kumar, M., Sahoo, D. K., Nain, N., Yadav, M., & Dadhich, V. (2025). Mapping the landscape of women entrepreneurs in micro and small enterprises: Trends, themes, and insights through systematic review and text mining. *F1000Research*, 14. <https://doi.org/10.12688/f1000research.163030.1>
 57. Narayanan, A., Mathur, A., Chetty, M., & Narayanan, A. (2020). Dark patterns: Past, present, and future. Retrieved from <https://darkpatterns.org>
 58. Nouwens, M., Liccardi, I., Veale, M., Karger, D., & Kagal, L. (2020). Dark patterns after the GDPR: Scraping consent pop-ups and demonstrating their influence. *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*, 1–13. <https://doi.org/10.1145/3313831.3376321>
 59. Pasquale, F. (2015). *The black box society: The secret algorithms that control money and information*. Harvard University Press.
 60. Petty, R. E., & Cacioppo, J. T. (1986). The elaboration likelihood model of persuasion. *Advances in Experimental Social Psychology*, 19, 123–205.
 61. Pankaj, Yadav, R., Naim, I., Kumar, M., Misra, S., Dewasiri, N. J., Rathnasiri, M. S. H., Yadav, M., Balodi, A., Kar, S., & Vihari, N. S. (2023). An Exploratory Study to Identify Effects of Blockchain Technology on Organizational Change and Practices. *2023 IEEE Technology & Engineering Management Conference - Asia Pacific (TEMSCON-ASPAC)*, 1–8. <https://doi.org/10.1109/TEMSCON-ASPAC59527.2023.10531372>
 62. Raji, I. D., Smart, A., White, R. N., Mitchell, M., Gebru, T., Hutchinson, B., ... & Barnes, P. (2020). Closing the AI accountability gap: Defining an end-to-end framework for internal algorithmic auditing. *Proceedings of the 2020 Conference on Fairness, Accountability, and Transparency*, 33–44. <https://doi.org/10.1145/3351095.3372873>
 63. Sandvig, C., Hamilton, K., Karahalios, K., & Langbort, C. (2014). Auditing algorithms: Research methods for detecting discrimination on internet platforms. In *Data and Discrimination: Collected Essays* (pp. 6–10).
 64. Shin, D. (2021). User perceptions of algorithmic decisions in the personalized AI system: Perceptual evaluation of fairness, accountability, transparency, and explainability. *Journal of Broadcasting & Electronic Media*, 65(1), 24–50. <https://doi.org/10.1080/08838151.2020.1851681>
 65. Shneiderman, B. (2020). Human-centered artificial intelligence: Reliable, safe & trustworthy. *International Journal of Human-Computer Interaction*, 36(6), 495–504. <https://doi.org/10.1080/10447318.2020.1741118>
 66. Sinha, R., & Singh, R. (2022). Building responsible AI: The role of interface clarity and user-centricity in digital persuasion. *AI & Society*. <https://doi.org/10.1007/s00146-022-01414-w>
 67. Sunstein, C. R. (2015). *Choosing not to choose: Understanding the value of choice*. Oxford University Press.
 68. Sushma, Malhan, D., & Kumar, M. (2025). Exploring Workplace Dynamics: A Systematic Review and Bibliometric Analysis of Work Stress, Gender Diversity, and Job Satisfaction in Higher Education Institutions. *Journal of Information Systems Engineering and Management*, 10(5(s)), 786 – 803. <https://doi.org/10.52783/jisem.v10i5s.772>
 69. Tax, S. S., Brown, S. W., & Chandrashekar, M. (1998). Customer evaluations of service complaint experiences: Implications for relationship marketing. *Journal of Marketing*, 62(2), 60–76.
 70. Utz, C., Specht, F., & Degeling, M. (2022). How cookie consent interfaces design affects user choices and autonomy. *Information Systems Journal*, 32(2), 261–287. <https://doi.org/10.1111/isj.12328>
 71. Voigt, P., & Von dem Bussche, A. (2017). *The EU General Data Protection Regulation (GDPR): A practical guide*. Springer.
 72. 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/ixp005>
 73. Wachter, S., Mittelstadt, B., & Floridi, L. (2021). 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.
 74. Wirtz, J., Kunz, W. H., Paluch, S., & Martins, A. (2022). Responsible AI in service: A research agenda. *Journal of Service Management*, 33(3), 441–456. <https://doi.org/10.1108/JOSM-10-2021-0364>

75. Wirtz, J., Patterson, P. G., Kunz, W. H., Gruber, T., Lu, V. N., Paluch, S., & Martins, A. (2022). Brave new world: Service robots in the frontline. *Journal of Service Management*, 33(4), 633–656. <https://doi.org/10.1108/JOSM-04-2022-0129>
76. Yeung, K. (2017). Hypernudge: Big data as a mode of regulation by design. *Information, Communication & Society*, 20(1), 118–136.
77. Yao, Y., Kim, J., & Kim, H. (2021). Designing to empower: Investigating the effects of transparency and control on user trust and autonomy in recommender systems. *Proceedings of the ACM on Human-Computer Interaction*, 5(CSCW2), 1–25. <https://doi.org/10.1145/3479543>
78. Yadav, M., Mittal, S., Kumar, M., Sahoo, A., Jayarathne, P. G. S. A., Yadav, M., Mittal, S., Kumar, M., Sahoo, A., & Jayarathne, P. G. S. A. (2024). From Textbooks to Teamwork. In <https://services.igi-global.com/resolvedoi/resolve.aspx?doi=10.4018/979-8-3693-3856-8.ch011> (pp. 333–352). IGI Global Scientific Publishing. <https://doi.org/10.4018/979-8-3693-3856-8.ch011>
79. Zuboff, S. (2019). *The age of surveillance capitalism: The fight for a human future at the new frontier of power*. PublicAffairs.

How to cite this article?

Magar, Dr. A., Kad, A., Kotwal, S., & Kulkarni, R. (2025). Taxing the Dream Home: A Critical Analysis of GST's Influence on Real Estate Affordability. *International Journal of Advance Research in Computer Science and Management Studies*, 13(6), 15–22 <https://doi.org/10.61161/ijarcsms.v13i6.3>