

Research-Thesis: Ethical customisation in AI travel booking: a between-subjects experiment on perceived agency, trust, and satisfaction

Auteur : Gennen, Jana

Promoteur(s) : El Midaoui, Youssra

Faculté : HEC-Ecole de gestion de l'Université de Liège

Diplôme : Master en sciences de gestion, à finalité spécialisée en international strategic marketing

Année académique : 2025-2026

URI/URL : <http://hdl.handle.net/2268.2/25825>

Avertissement à l'attention des usagers :

Tous les documents placés en accès ouvert sur le site le site MatheO sont protégés par le droit d'auteur. Conformément aux principes énoncés par la "Budapest Open Access Initiative"(BOAI, 2002), l'utilisateur du site peut lire, télécharger, copier, transmettre, imprimer, chercher ou faire un lien vers le texte intégral de ces documents, les disséquer pour les indexer, s'en servir de données pour un logiciel, ou s'en servir à toute autre fin légale (ou prévue par la réglementation relative au droit d'auteur). Toute utilisation du document à des fins commerciales est strictement interdite.

Par ailleurs, l'utilisateur s'engage à respecter les droits moraux de l'auteur, principalement le droit à l'intégrité de l'oeuvre et le droit de paternité et ce dans toute utilisation que l'utilisateur entreprend. Ainsi, à titre d'exemple, lorsqu'il reproduira un document par extrait ou dans son intégralité, l'utilisateur citera de manière complète les sources telles que mentionnées ci-dessus. Toute utilisation non explicitement autorisée ci-avant (telle que par exemple, la modification du document ou son résumé) nécessite l'autorisation préalable et expresse des auteurs ou de leurs ayants droit.

**ETHICAL CUSTOMISATION IN AI TRAVEL
BOOKING: A BETWEEN-SUBJECTS
EXPERIMENT ON PERCEIVED AGENCY, TRUST,
AND SATISFACTION**

Jury:
Supervisor:
Youssra EL MIDAOUI
Reader(s):
Nadia STEILS

Master thesis presented by
Jana GENNEN
To obtain the degree of
MASTER IN MANAGEMENT
with a specialization in
INTERNATIONAL STRATEGIC
MARKETING
Academic year 2025/2026



Acknowledgements

I would like to begin by expressing my sincere gratitude to my supervisor, Youssra El Midaoui, for her guidance throughout this thesis. I deeply appreciate the time she dedicated to me, the thoughtful guidance she shared, and her willingness to provide insightful answers to my many questions. Her advice and encouragement have been extremely valuable and helped me move forward, even during more challenging moments.

I am equally grateful to my family, whose constant encouragement and belief in me helped me enormously through not only this thesis, but my whole academic journey. Their support and guidance kept me grounded during every challenging moment.

I also want to thank my friends for their help and feedback, as well as for their support and motivation during stressful periods. Their understanding and encouragement made completing this thesis much easier than it would have been otherwise.

Finally, I would like to express great appreciation to all participants who generously took the time to take part in the survey. Their contribution was essential to the successful completion of this thesis.

Table of Contents

Chapter 1: Introduction	2
1.1 Context	2
1.2 Research motivations	3
1.2.1 Managerial motivations	3
1.2.2 Academic motivations	3
1.3 Problem statement	3
1.4 Research approach and theoretical framework	4
1.5 Contributions	4
Chapter 2: Literature review	5
2.1 Artificial Intelligence	5
2.1.1 Definition of Artificial Intelligence	5
2.1.2 Artificial intelligence technologies	5
2.1.2.1 <i>Machine learning</i>	5
2.1.2.2 <i>Deep Learning and Neural Networks</i>	5
2.1.2.3 <i>Natural Language Processing and Large Language Models</i>	6
2.2 Recommender Systems in the Tourism Ecosystem	6
2.2.1 Definition	6
2.2.2 Mechanisms and algorithms	6
2.3 Ethical AI	7
2.3.1 Ethical paradigms in AI	7
2.3.2 Algorithmic fairness and the mechanics of transparency	8
2.3.3 Human agency	8
2.3.4 Perceived ethical alignment	9
2.3.5 Customisation	10
2.3.6 Signalling Theory as a Complementary Lens	10
2.4 Traveller satisfaction	11
2.5 Research model and hypotheses development	12
2.5.1 Conceptual research model	12
2.5.2 Hypotheses development	12
2.5.2.1 <i>Ethical customisation</i>	12
2.5.2.2 <i>Mediator: Perceived agency</i>	13
2.5.2.3 <i>Mediator: Platform trust</i>	13
2.5.2.4 <i>Mediator: Perceived ethical alignment</i>	13
2.5.2.5 <i>Moderator: Algorithmic literacy</i>	14
2.5.2.6 <i>Control variable: perceived price</i>	14

2.5.3	Conceptual model	15
Chapter 3: Research design.....		16
3.1	Methodology.....	16
3.2	Sample	16
3.3	Research instrument.....	17
3.4	Pilot test	20
Chapter 4: Results		22
4.1	Preliminary Analysis: Scale Reliability and Structural Validity.....	22
4.1.1	Internal Consistency Reliability	22
4.1.2	Structural validity: Exploratory Factor Analysis	23
4.1.3	Supplementary Confirmatory Evidence: Single-factor CFA.....	24
4.2	Sample description and Randomisation Check.....	25
4.3	Descriptive Statistics and Correlations	27
4.4	Hypothesis Testing.....	28
4.4.1	H1: Direct Effect of Ethical Customisation on Satisfaction	28
4.4.2	H2: Mediation by Perceived Agency	30
4.4.3	H3: Mediation by Platform Trust.....	30
4.4.4	H4: Mediation by Ethical Alignment	31
4.4.5	H5: Moderation by Algorithmic Literacy.....	31
4.4.6	Summary of Hypothesis Testing	33
4.5	Exploratory Analysis: Behavioural Intention	34
4.6	Qualitative Comments: Supplementary Insights	34
Chapter 5: Discussion		36
5.1	The Absence of a Direct Effect: H1	36
5.2	Perceived Agency as a Partial Mediator: H2	36
5.3	Platform Trust as a Non-Mediator: H3.....	37
5.4	Perceived Ethical Alignment as a Non-Mediator: H4	37
5.5	The Role of Algorithmic Literacy: H5.....	38
5.6	Synthesis and Theoretical Contributions	39
Chapter 6: Conclusions		41
6.1	Summary	41
6.2	Managerial implications	41
6.3	Theoretical implications	42
6.4	Limitations and implications for future research	42
Appendices		44
Appendix A: Survey Instrument (English Version)		44

Appendix B: Experimental Interface Stimuli	48
Appendix C: Supplementary Statistical Output	50
C1 Reliability statistics for all scales (Cronbach's α , item-dropped statistics).....	50
C2 Exploratory Factor Analysis factor loading table.....	53
C3 Single-Factor Confirmatory Factor Analyses	54
C4 Full descriptive statistics output, including distribution plots.....	58
C5 Frequency tables for all constructs	59
C6 Correlations.....	61
C7 Assumption Checks	62
C8 Hypothesis Testing	63
C9 Exploratory Analysis	73
List of Resource Persons	74
References	75

List of Abbreviations

Abbreviations	Signification
AI	Artificial Intelligence
CFA	Confirmatory Factor Analysis
DL	Deep Learning
EDP	Expectancy-Disconfirmation Paradigm
EFA	Exploratory Factor Analysis
EVM	Experimental Vignette Methodology
HITL	Human-In-The-Loop
HOOL	Human-Out-Of-Loop
HOTL	Human-On-The-Loop
LLM	Large Language Model
ML	Machine Learning
NLP	Natural Language Processing
RS	Recommender System
S-O-R	Stimulus-Organism-Response
VSD	Value-Sensitive Design
WLSMV	Weighted Least Squares Mean and Variance
XAI	Explainable AI

Chapter 1: Introduction

1.1 Context

Driven by advancements in computing and technology, the concept of developing systems that can carry out tasks typically associated with human intelligence has evolved gradually over time. Artificial Intelligence (AI) as a formal research field, however, emerged from the 1956 Dartmouth workshop organised by J. McCarthy, which marked the beginning of scientific efforts to model human reasoning through computational systems. A further advance in this field occurred in 1958, when Frank Rosenblatt introduced the concept of the “perceptron”. As one of the earliest learning-based systems, the perceptron demonstrated how machines could adapt to environmental patterns without explicit rule-based programming (Fanti et al., 2022).

In recent years, AI applications have expanded beyond controlled environments and become increasingly embedded within the tourism industry. Since tourism is a crucial component of global socioeconomic development, it has taken the lead in striving for technological optimisation to promote international mobility (Chen & Wei, 2024). This transformation gave rise to “Smart Tourism”, where data-driven insights manage the challenges of travel logistics and traveller behaviour. Through the integration of digital data infrastructures within physical travel environments, “smart business ecosystems” are created, allowing real-time interactions between travellers and service providers (Gretzel et al., 2015). This includes the optimisation of business operations, predictive maintenance in hospitality, and the enhancement of customer relationship management (Cristian et al., 2024). As AI technologies continue to evolve beyond operational support toward more interactive and personalised applications, the transition from static recommendation engines to proactive AI assistants is exemplified by the 2024 launch of Expedia’s “Romie”, using Deep Learning (DL) to tailor travel experiences (Lim & Kim, 2025).

Since the digital travel landscape continues to develop, AI-powered Recommender Systems (RS) have become increasingly important in assisting with the vast amount of travel data available. By analysing individual interests and preferences, these systems provide personalised recommendations that simplify decision-making and enhance the overall travel planning experience (Sarkar et al., 2023). However, the rise of AI in travel also brings its downsides. While efficient, these systems often operate as “Black Boxes”, which are computational systems where the internal logic and decision-making processes remain hidden from the user. This lack of clarity makes it hard for users to understand why a certain result was generated, creating what researchers call a “transparency gap” (Kopalle et al., 2022). In travel, where decisions carry high financial and time costs, this opacity has the potential to fuel consumer anxiety and scepticism (Pavlou, 2003).

Besides creating user uncertainty, the lack of insight into how algorithms work also raises ethical concerns. One of these issues is algorithmic bias, where systems may unintentionally reproduce human prejudices. By systematically favouring partnered hotel chains over independent properties regardless of traveller fit, traveller trust and satisfaction can quickly weaken (I. Koo et al., 2025). To counter these risks, researchers have introduced frameworks such as “Trustworthy AI” and “Explainable AI” (XAI), which prioritise transparency and interpretability. This helps to ensure that technologies respect fundamental human rights, reduce user anxiety and foster long-term confidence (Santosh & Wall, 2022). In the travel industry, XAI serves as a bridge, translating complex data processing into human-readable justifications that strengthen the perceived reliability of the platform.

While existing platforms personalise recommendations through system-driven algorithms, this thesis examines a fundamentally different approach by allowing users to customise the ethical parameters of the AI themselves, which shifts moral agency from the platform to the traveller. Personalisation is typically a system-driven process, where algorithms predict user needs based on implicit data (Samala et al., 2022), while customisation is a user-driven process that allows travellers to have the control to

actively define their own parameters (Zhang & Sundar, 2019). In a digital world, where “algorithm fatigue” frequently leaves users feeling overloaded with automated recommendations that do not truly fit with who they are, regaining this sense of agency becomes a psychological necessity (Yang et al., 2024).

1.2 Research motivations

1.2.1 Managerial motivations

As the tourism industry increasingly relies on AI to manage the growing volume and complexity of global travel data, AI-powered RSs have appeared as an important mechanism for delivering personalised services (Zhao et al., 2024). These systems optimise operational processes and improve the relevance and the quality of information presented to the users. However, their internal functioning often remains opaque, which leaves users unaware of how their recommendations are generated. Even though AI technologies increase efficiency and enhance user experiences, their “Black Box” nature can intensify concerns around transparency, possibly leading to consumer unease and distrust (Singu et al., 2025). From a managerial perspective, issues such as algorithmic bias and the perception of manipulation of users toward specific travel options can result in significant reputational damage and a decline in consumer confidence (Milano et al., 2020). Research indicates that ethical transparency, the clear communication of AI practices and the assurance of algorithmic fairness are no longer just a regulatory requirement but a critical driver of brand loyalty and positive brand perception (Vasiliki et al., 2025). Consequently, this thesis supports tourism professionals by exploring how ethical customisation can serve as a competitive differentiator through the transformation of transactional interactions into collaborative partnerships that build long-term customer confidence.

1.2.2 Academic motivations

Beyond its practical value for managers, this thesis aims to address a clear gap in the academic literature at the intersection of AI ethics and user agency in digital marketplaces. Although existing studies have extensively explored the functional benefits of AI personalisation, such as service speed and accuracy, the psychological impact of moral self-determination remains under-researched (Lu, 2024). In this study, ethical customisation is positioned as an essential stimulus that affects both the user’s cognitive evaluations and emotional reactions, building on and extending the Stimulus-Organism-Response (S-O-R) framework in an AI-powered tourism context (Park & Son, 2025). While scholars have identified that transparency significantly strengthens trust through enhanced perceptions of fairness, there is a lack of empirical work that examines how granting users direct control over ethical attributes affects satisfaction (Milano et al., 2020). Furthermore, much of the tourism literature often lacks a focus on the “Human-in-the-loop” (HITL) philosophy (Raees et al., 2024). By examining the role of user-driven ethical customisation, this study investigates how perceived agency helps overcome algorithm aversion, ultimately contributing to the advancement of more human-centric and responsible AI systems in tourism.

1.3 Problem statement

While current platforms excel at personalising for functional preferences like price or location, there is an emerging demand for the customisation of ethical attributes. As travellers become increasingly aware of data protection and social impact, the ability to configure ethical preferences, such as how data is used or which criteria are prioritised, is becoming a fundamental component of the booking experience (I. Koo et al., 2025). This shift toward value co-creation suggests that satisfaction is no longer solely about the product, but about the alignment between the traveller’s moral values and the platform’s recommendation logic (Ali et al., 2016). When a platform offers these ethical controls, it signals a commitment to the user’s personal integrity, transforming the booking experience from a one-sided data extraction into a co-created exchange of values. However, the “Black Box” nature of most booking platforms denies travellers the ability to align their bookings with these values (I. Koo et al., 2025).

This lack of active human decision-making processes can result in cognitive dissonance, where a traveller may feel satisfied with the price but ethically compromised by the platform's lack of transparency. If travellers cannot customise the ethical logic of the AI, the resulting recommendations can be perceived as manipulations or nudges rather than value-adding services (Kopalle et al., 2022). By aligning AI ethics, user agency, and customer psychology, this thesis aims to address whether granting users the agency to explicitly customise ethical preferences can serve as a corrective measure for the trust deficit in AI-powered travel recommendations and ultimately enhance long-term satisfaction.

1.4 Research approach and theoretical framework

The research methodology used in this study is deductive, employing the S-O-R framework as a theoretical anchor to evaluate the user experience. According to this framework, external influences can shape travellers' thoughts and perceptions, which can ultimately affect the way they respond or behave (Park & Son, 2025). The Stimulus (S) is operationalised as the perceived availability of ethical customisation features within the AI-powered booking platform, allowing users to make decisions based on moral criteria. The Organism (O) represents the internal psychological reactions of the traveller, specifically their perceived sense of agency, their trust levels, and their perception of ethical alignment (Pavlou, 2003). Finally, the Response (R) is measured through traveller satisfaction and ultimately the behavioural intention to return to the platform.

By implementing a quantitative methodology, this thesis aims to validate the causal link between user-driven ethical control and positive user outcomes. This perspective is consistent with the HITL design philosophy, which suggests that optimal AI performance requires human ethical judgement to remain central, by being complemented rather than displaced by AI (Santosh & Wall, 2022). This aligns with Lim & Kim's (2025) notion of consumer empowerment, where users retain autonomy rather than simply opting out of personalisation. Their empirical findings suggest that empowerment acts as a "protective mechanism", particularly for users who may lack deep technical knowledge of AI (Lim & Kim, 2025). However, while previous research has viewed empowerment primarily as the ability to opt out of personalisation, this thesis argues for a more specific approach of allowing users to customise the ethical parameters themselves to ensure the AI's logic matches their individual moral standards.

1.5 Contributions

This thesis aims to make four significant contributions to the literature. The most concrete is experimental by treating an ethical customisation interface as the stimulus in the S-O-R framework (Mehrabian & Russell, 1974). This provides causal evidence for a mechanism that prior work had only theorised. The second contribution builds directly on Zhang & Sundar's (2019) personalisation-customisation distinction. This thesis tests ethical parameter control, which is a quantitatively different form of user empowerment, and one with stronger moral implications. Third, perceived ethical alignment is introduced here as a construct capturing real-time value resonance at first interface exposure, while the null mediation result clarifies, rather than undermines, the conditions under which it operates. Finally, the consistent main effect of algorithmic literacy across all model specifications establishes it as a robust individual-level predictor of AI-driven experience, independently of experimental condition. This is a pattern with direct implications for segmented interface design in AI-powered tourism platforms.

Chapter 2: Literature review

2.1 Artificial Intelligence

2.1.1 Definition of Artificial Intelligence

Although Machine (Artificial) Intelligence has been defined and interpreted in many different ways, it is formally characterised as the practical, real-world ability of artificial or non-human systems to execute tasks, engage in communication and interactions, solve problems, and apply logical reasoning in a manner similar to that of humans (Gil De Zúñiga et al., 2024). AI is therefore defined as the capability of a system to accurately analyse external information, learn from that information, and apply the required knowledge to accomplish particular objectives and tasks through adaptive and flexible responses (Haenlein & Kaplan, 2019). By mimicking human cognitive capabilities, AI resolves complex computational problems far faster than humans by applying algorithmic rules to analyse and process the external data (Ferrer-Rosell et al., 2023).

2.1.2 Artificial intelligence technologies

2.1.2.1 *Machine learning*

Machine learning (ML) serves as one of the foundational engines of AI, referring to algorithms that analyse data patterns and human-machine interactions to interpret and act upon data autonomously (Kopalle et al., 2022). Rather than requiring explicit programming, ML systems improve iteratively as they process more data. A key technical challenge is preventing models from becoming too finely calibrated to their training examples at the expense of broader applicability, also referred to as “overfitting” (Janiesch et al., 2021).

Within tourism RSs, ML is the mechanism through which platforms like Booking.com translate accumulated user behaviour into personalised hotel suggestions. By collecting personal data history, navigation patterns, and past booking interactions, these systems create a self-reinforcing cycle, meaning more data yields more precise recommendations, and more precise recommendations drive deeper user engagement (Kopalle et al., 2022; M. Li et al., 2021). The ultimate objective is personalisation, referring to the automated adaptation of outputs to individual user preferences (Samala et al., 2022). Yet this very strength gives rise to the central tension of this thesis. Because ML models optimise for prediction accuracy rather than interpretability, their internal logic remains opaque for the user. The more data a system accumulates and the more precisely it personalises, the less the traveller understands why a particular hotel was recommended, laying the groundwork for what later sections identify as the “transparency gap” (Kopalle et al., 2022).

2.1.2.2 *Deep Learning and Neural Networks*

Building on ML, DL relies on neural networks with multiple layers to automatically detect complex patterns within large and unstructured datasets, eliminating the need for specifically programmed rules (Janiesch et al., 2021). These systems, which are broadly modelled after neurological structures, are made up of nodes that resemble neurons linked together in ways that are gradually adjusted during the training process (Kopalle et al., 2022). In the context of tourism RSs, DL enables a significant improvement in recommendation quality by automatically identifying subtle regularities in traveller behaviour, which outperforms traditional rule-based approaches (Janiesch et al., 2021).

However, this gain in performance comes at a direct cost to transparency. The multi-layered architecture that makes DL so powerful is precisely what makes its outputs so difficult to understand. The reasoning behind any individual recommendation is distributed across hundreds or thousands of weighted connections, none of which are legible to the user. As a result, the more sophisticated the DL-powered RS, the more the traveller experiences its outputs as arbitrary or unaccountable.

2.1.2.3 *Natural Language Processing and Large Language Models*

Natural Language Processing (NLP) enables AI systems to read and make sense of human language, with transformer-based architectures making it possible for models to grasp the relational and contextual dimensions of language rather than treating words as individual components (Janiesch et al., 2021). These sophisticated frameworks are the foundation for RSs, especially in their ability to extract rich insights from user-generated content like customer reviews (Kopalle et al., 2022). The integration of DL-based NLP allows RSs to perform holistic learning, helping the systems to understand the “human” element of data.

Large Language Models (LLMs) represent the shift in the digital tourism landscape by enabling the transition from static information retrieval to dynamic, human-centric interaction (Gu, 2024). Unlike traditional RSs that rely on rigid, pre-programmed filters, LLMs utilise their computational depth to comprehend complex linguistic nuances and implicit traveller preferences (González-Sanz et al., 2025). In the tourism sector, LLMs facilitate “hyper-personalisation” by acting as proactive travel assistants that can interpret the multifaceted value preferences of travellers expressed in natural language (Bramer & Stahl, 2025).

This is particularly relevant to ethical customisation, as LLMs can provide real-time, human-readable justifications for proposed recommendations, thereby bridging the transparency gap (Samala et al., 2022). By synthesising vast amounts of unstructured data ranging from travellers’ blogs to real-time reviews, LLMs empower the HITL design philosophy, where AI enhances rather than replaces human moral judgment (Santosh & Wall, 2022). However, the increased complexity of LLMs further intensifies the “Black Box” challenge, necessitating the implementation of XAI methods to maintain user trust and agency (McDermid et al., 2021).

2.2 Recommender Systems in the Tourism Ecosystem

2.2.1 Definition

Tourism is a dynamic global industry that contributes significantly to the global economy by creating opportunities for economic growth and supporting the development of regions around the world (Kontogianni et al., 2024). The UNWTO defines tourism as a complex activity that carries social and cultural significance and involves people temporarily travelling outside their usual environment (Gretzel et al., 2015). The sector’s global scale and the volume of traveller decisions it processes each year make it one of the most data-intensive industries in the world, and a natural domain for AI-powered optimisation (Kontogianni et al., 2024). Contemporary Smart Tourism frameworks integrate digital data infrastructures with physical travel environments, creating ecosystems where AI-powered RSs serve as the primary mechanisms for anticipating and responding to traveller needs (Gretzel et al., 2015). By employing neural networks and LLMs, these systems have evolved from static information retrieval tools into proactive assistants capable of human-like contextual responses (Kontogianni et al., 2024). Contemporary RSs integrate behavioural, contextual, and real-time data to generate increasingly personalised and adaptive travel recommendations (Walek & Sládek, 2025).

2.2.2 Mechanisms and algorithms

In the tourism industry, AI-powered systems aim to enhance recommendation accuracy by gathering rich user features from multisource data, employing Deep Neural Networks (DNN) to predict user preferences more precisely (Lin et al., 2025). In the context of booking platforms, these systems have the power to improve not only the accuracy but also the user’s overall experience, leading to stronger engagement and promotional outcomes (Solano-Barliza et al., 2024). The integration of AI has already improved travel website performance by reducing the computational resources required while maintaining service quality (Begu et al., 2025). The AI-powered RSs fall under the classification of Limited AI, performing predefined tasks, such as providing personalised recommendations and the use of smart chatbots (Milwood et al., 2023). Understanding how AI-powered recommendation

systems function within the tourism industry requires examining the technologies and processes that support them. Similar to other sectors, tourism platforms use ML methods and optimisation techniques to interpret and organise the vast data collected from users.

Solano-Barliza et al. (2024) identify several dominant typologies of RSs applied to tourism. The tourism industry standard for achieving high performance is the hybrid model (Solano-Barliza et al., 2024), which creates a “virtuous cycle” by strategically combining two or more techniques to optimise the recommendation accuracy. This approach combines the strengths of different recommendation techniques, including collaborative filtering, content-based filtering, context-aware filtering, and demographic-based methods. Among these, collaborative filtering is the most used technique in the hybrid system, as it enables the platform to evaluate user ratings and identify similarities between users with similar preferences. Context-aware systems compile contextual variables related to the user’s activity, such as location, time of the day and weather (Solano-Barliza et al., 2024). The recent integration of LLMs has further enhanced this cycle. While traditional neural models can find it difficult to fully capture the complexity of text, LLMs use advanced semantic understanding to produce predictions that resemble human-like reasoning (Zhao et al., 2024).

However, a main psychological challenge, the “Black Box” effect, remains. By relying on hidden algorithmic layers, these systems often lead to a trust deficit, where the logic behind a recommendation is hidden from the user, making it difficult to recommend diverse or exploratory items. This persistent challenge suggests that technical precision is no longer the main determinant of traveller satisfaction. As systems become more autonomous, a “transparency gap” emerges between algorithmic efficiency and user understanding of the system (Solano-Barliza et al., 2024). Consequently, a critical paradox emerges: an inverse relationship often exists between algorithmic optimisation and the traveller’s perceived autonomy, as high-precision opaque models frequently erode the user’s sense of agency. Moving beyond the “Black Box” requires a transition toward ethical AI principles that prioritise user agency. By shifting from purely automated processing to collaborative interaction models, the industry can restore the human element in digital booking. LLMs provide thus a unique opportunity, as their reasoning abilities can be utilised to offer explainable recommendations (Zhao et al., 2024). This theoretical shift toward user-centric controls, such as explicit customisation, forms the basis for the following exploration of perceived agency, ethical alignment and trust within the traveller’s digital journey.

2.3 Ethical AI

2.3.1 [Ethical paradigms in AI](#)

While AI-powered recommendation systems are increasingly necessary for navigating contemporary data saturation, their deployment introduces significant ethical vulnerabilities, specifically regarding data privacy, identity security, and algorithmic deception (Bhardwaj et al., 2025). In response to these developments, the domain of Ethical AI has emerged, relying on moral frameworks that guide the design, deployment, and governance of automated systems (Woodgate & Ajmeri, 2024). As AI systems progressively undertake functions that demand human-level cognitive capabilities, the principles of Trustworthy AI and XAI have developed as central frameworks aimed at promoting operational transparency and protecting fundamental human rights (Santosh & Wall, 2022). To evaluate trust within these architectures, one must distinguish between functionality trust, which focuses on technical accountability and robustness, and human-like trust, which evaluates the system’s perception (also including societal and environmental care) (Choung et al., 2023). Ethical assessment includes psychological (cognitive liberty, transparency), social (fairness), and political (institutional stability) dimensions (Kazim & Koshiyama, 2021). Ultimately, the integration of values such as impartiality and accountability drives user acceptance, particularly in smart tourism, where increased awareness of data protection makes ethical standards a prerequisite for adoption (I. Koo et al., 2025).

2.3.2 [Algorithmic fairness and the mechanics of transparency](#)

A primary ethical requirement for AI-powered RSs is the pursuit of fairness by delivering equitable outcomes that avoid discriminatory treatment. This concept is rooted in human equality and encompasses various theories (including corrective and distributive fairness), which can be evaluated through the lenses of the treatment process or the ultimate impact. Central to this pursuit is the mitigation of bias, representing discriminatory patterns that are often inherited from historical datasets. These patterns can lead to the purposeful exploitation of marginalised groups or differences in service quality. Accessibility should therefore be a main goal to guarantee that such systems continue to be inclusive and fair. Participatory governance can help ensure that AI tools are both reasonably priced and flexible enough to meet the needs of a wide range of users (Kazim & Koshiyama, 2021). As ML increasingly controls sectors critical to human well-being, the risk of bias and opacity threatens to erode trust (McDermid et al., 2021). Mounting evidence indicates that hiring and credit-related systems have displayed gender- and race-based biases, while additional misinterpretations can be linked to shortcomings in system design (Kazim & Koshiyama, 2021). Nowadays, algorithmic transparency is a fundamental ethical dimension of AI, making it a key requirement for building trustworthy AI systems. Transparency is closely linked with accountability, encouraging users to share data by reinterpreting the decision-making process (Choung et al., 2023).

To bridge the gaps between trust and satisfaction, booking platforms are increasingly adopting XAI methods to identify and mitigate bias-inducing models. If users perceive the systems and their mechanisms as incomprehensible, trust might not automatically translate into user satisfaction. However, as these AI-powered RSs provide transparency, the relationship between trust and satisfaction might be strengthened (Hassan et al., 2025). On one hand, these XAI methods can be categorised by their scope, describing local explanations for individual predictions versus global explanations for general model behaviour. On the other hand, they can be categorised by their timing, distinguishing between contemporaneous insights (will be available at the same time as the predictions) and post-hoc explanations (which will be provided after the prediction is made). A post-hoc local explanation would be necessary in case an individual wants to challenge an automated prediction made about them under data protection law, for example. By utilising techniques like feature importance methods or DeepLIFT, developers can reveal which data inputs carry the most weight in the decision-making process of an algorithm, thereby enabling users to provide truly informed consent to recommendations made by the ML-based systems (McDermid et al., 2021).

2.3.3 [Human agency](#)

The ethical principles of freedom and autonomy, which together constitute the notion of human agency, serve as the foundation for the development and use of AI (Choung et al., 2023). In a human-centric approach to AI, maintaining agency is a prerequisite for reassuring human dignity, referring to the recognition of the moral status of individuals as rational agents capable of autonomous existence. To analyse this framework, a distinction must be made between human agency, namely the capacity for “moral self-determination”, and control, which refers to the power to direct the system’s outputs. Support for this perspective can be found in the self-determination theory, which positions autonomy as a fundamental condition for sustaining intrinsic motivation and influencing how systems are evaluated (Ryan & Deci, 2020). Within this philosophical context, autonomy can be further divided into mental autonomy, which safeguards a person’s deliberative abilities and internal thought processes, and physical autonomy, which ensures respect for a person’s body and their external choices. AI-powered RSs protect user independence by building in features for informed consent and manual control, while still relying on human experts to monitor the outcomes (Kazim & Koshiyama, 2021). Enabling users to maintain their accountability by giving their meaningful and informed consent while interacting with an AI-powered RS is essential to protect them from manipulation, all while retaining the right to revoke consent.

Scholars distinguish between varying degrees of human participation in AI governance. Human-in-the-Loop (HITL) describes a state where humans intervene in specific decisions; Human-on-the-Loop (HOTL), where humans monitor the overall system; and Human-out-of-the-Loop (HOOL), which refers to a completely independent state (Santosh & Wall, 2022).

Since the decision-making process has shifted from humans to machines when interacting with AI-powered RSs, the provision of transparent explanations becomes vital for maintaining user control and fostering justified confidence (McDermid et al., 2021). Scholars have expressed concern that the widespread deployment of AI systems may unintentionally weaken people's capacities for reasoning and reflection, potentially reducing deliberative agency (Kazim & Koshiyama, 2021). Human agency in digital travel environments is not merely the ability to select a button, but the compound experience of initiation, control, and the capacity for valid decision-making. Ethical customisation features give travellers a way to make sure their personal beliefs and moral values are actually driving the decisions they make on a platform (Pacherie, 2007). When the system visibly responds to ethical preferences, it reframes the interaction where the traveller is no longer the object of personalisation but its author (Siegel, 2025).

2.3.4 Perceived ethical alignment

Perceived ethical alignment is a concept that has the ability to connect user autonomy with trust in digital platforms. It describes the extent to which individuals believe that the recommendations provided by an AI system reflect and respect their own moral values and ethical expectations. Unlike platform trust, which relates to beliefs about a system's reliability and benevolence (Jiang et al., 2025), and unlike perceived agency, which refers to the user's sense of control over system outcomes (Ryan & Deci, 2020), perceived ethical alignment focuses specifically on value coherence (whether the AI appears to reflect, respect, or align with the traveller's ethical identity).

This construct is supported by two complementary streams of literature. First, Koo et al. (2025) show that the adoption of AI-powered smart tourism technologies is influenced not only by functional performance but also by the perceptions that systems respect ethical considerations and personal preferences. This suggests that users evaluate AI systems through a moral as well as a functional lens. Second, the Value-Sensitive Design (VSD) literature (Del Valle & Lara, 2024) argues that when technologies align with users' moral values, they produce distinct evaluative responses that go beyond traditional satisfaction or usability judgements. In this context, when a recommendation system prioritises options such as eco-certified or locally owned hotels, users may experience something more than convenience. They may feel that the system is creating a sense of moral recognition that strengthens their overall evaluation of the platform.

Unlike AI ethics perception scales (I. Koo et al., 2025), which measure users' beliefs about whether a system follows ethical principles in the abstract, perceived ethical alignment measures a personalised, real-time experience of value coherence (whether this system, at this moment, appears to reflect *my* moral priorities). It is theoretically distinct from both trust and agency, since a user can trust a system and feel in control of it without experiencing that its outputs substantively reflect their moral values. Equally, a user may perceive strong value alignment in a system's outputs without trusting that the platform's underlying motives are genuinely user-centric rather than commercially driven. This is theoretically supported by VSD literature (Del Valle & Lara, 2024), which demonstrates that evaluative responses to a value-fit are affectively distinct from confidence-based or control-based evaluations.

Its inclusion in this study's model reflects a central contribution of this study. Ethical customisation, therefore, not only shapes user experience by increasing perceived control or confidence in the system, but also by creating a direct sense of moral resonance between user values and system outcomes. Together, these three pathways offer a more complete explanation of how ethical interface design influences user experience than approaches that consider trust or agency in isolation.

2.3.5 Customisation

Recent empirical research suggests that giving users control over the personalisation of AI recommendations significantly influences both user satisfaction and continued commitment to a brand (S. Y. Kim & Kim, 2025). By granting users granular control over ethical parameters, such as the prioritisation of specific ethical criteria or data usage preferences, booking platforms can significantly enhance trust (I. Koo et al., 2025). As users grow increasingly conscious of ethical standards, the ability to configure these preferences is no longer just a feature of convenience but a critical component of satisfaction in AI-integrated travel environments (I. Koo et al., 2025). This relationship can be effectively interpreted using the S-O-R framework first proposed by Mehrabian & Russell (1974). Within this model, the perceived availability of ethical customisation features functions as the Stimulus (S), prompting changes within the traveller as the Organism (O), particularly in relation to their internal perceived sense of agency (PSoA) (Park & Son, 2025). Unlike system-driven personalisation, which relies on implicit data, customisation is a user-driven process that empowers the traveller to proactively define their parameters. This represents a psychological necessity to combat “algorithmic fatigue” and ensure that recommendations reflect a user’s refined identity (Yang et al., 2024).

Zhang and Sundar’s (2019) distinction between system-driven personalisation and user-driven customisation has so far been applied mainly to functional features such as content filtering, price settings, or interface preferences. What has not yet been fully explored is whether this same distinction remains when the attributes in question are ethical rather than functional. It is still unclear whether giving users control over moral dimensions (such as sustainability preferences, data privacy levels, or social impact weighting) leads to different psychological outcomes compared to traditional forms of customisation. This thesis addresses that gap by empirically testing this extension. If users respond differently to ethical customisation rather than to functional customisation, it suggests that the personalisation-customisation framework is not equally distributed across all types of attributes. Instead, it may be that ethical control allows a deeper layer of user agency. In this sense, moral self-determination could represent a qualitatively distinct form of interaction with AI systems, with important implications for trust, satisfaction, and the design of digital platforms.

While Kim & Kim (2025) state that customisation increases satisfaction, only a few studies lay out the depth of control that the user desires or wants to tolerate. This suggests that, for a platform to achieve full user alignment, customisation must be implemented as a choice design that supports rather than overwhelms the user’s moral agency.

2.3.6 Signalling Theory as a Complementary Lens

A secondary theoretical perspective that contextualises the experimental design of this study is Signalling Theory (Spence, 1973), applied here not as a core explanatory framework but as a mechanism that accounts for why static interface cues alone may produce psychological responses. In situations where users cannot directly observe the true quality or intentions of a system, visible design features can act as signals that communicate the platform’s underlying values before any interaction even takes place. Applied to ethical customisation in booking platforms, features such as eco-friendliness or data privacy sliders can signal a commitment to transparency and values-aligned services, even when users do not actively engage with them. What matters is not only what these features allow users to do, but what their simple presence communicates. This perspective is particularly relevant to the experimental design of this study, which focuses on the perceptual effects of ethical customisation availability rather than its active use. In this context, the signal itself functions as the stimulus. This helps to explain why simply exposing users to an interface with ethical customisation features can generate meaningful psychological responses. Signalling Theory, therefore, provides a useful bridge between the XAI transparency literature and this thesis’s empirical approach by offering a theoretical hypothesis for why static interface cues alone may shape user perceptions and responses.

2.4 Traveller satisfaction

The rapid integration of AI within the tourism industry has radically shifted the experience economy from delivering standard services to the staging of meaningful, personalised interactions adapted to the users' needs and expectations (Lv et al., 2024). In this digital landscape, customer satisfaction is no longer viewed merely as a reaction to a single transaction, but rather as a global evaluation of the entire relationship history between the traveller and the platform (Filieri et al., 2015). According to Ranjbaran et al. (2024), this concept refers to the extent to which consumers derive satisfaction from a firm's offering and overall interactions. AI-powered travel recommendations serve as a significant driver for this satisfaction by generating proposals that align precisely with specific travellers' needs and interests (Hlee et al., 2025).

The psychological foundation of satisfaction is deeply rooted in the Expectancy-Disconfirmation Paradigm (EDP), which hypothesises that satisfaction or dissatisfaction arises from a process where the users weigh their expectations from the pre-use of a service to their post-use perceptions. When a service exceeds prior expectations, users experience a positive evaluative response, whereas when it falls short, disappointment follows. However, when expectations are precisely met, satisfaction remains neutral. This expectation-gap logic establishes how travellers evaluate AI-generated recommendations (Oliveri et al., 2019). As modern tourism transitions from product-centric to experience-centric models, the traveller's ability to actively co-create their journey through personalised interactions becomes the essential indicator of value (Ali et al., 2016).

However, a critical academic distinction must be drawn between the concepts "value" and "values". While traditional "value" refers to the functional utility derived from a single transaction based on preferences, such as price and location, "values" represent the broader social behaviours (including attitudes and ideology) of the user (Boksberger & Melsen, 2011). A prominent moral value in contemporary tourism, identified by Boksberger et al., is ethical action that benefits others, also known as morality.

When booking platforms integrate algorithms that allow for creative experiences, such as selecting destinations based on sustainability or community impact, traveller satisfaction increases in direct proportion to how effectively these systems facilitate these ethical purposes (Ali et al., 2016). This relationship, once again, can be operationalised through the S-O-R framework, which suggests that external stimuli, here the perceived availability of ethical customisation features such as sliders for sustainability or data privacy, can trigger a shift in internal organisms, specifically their cognitive sense of agency, their perception of ethical alignment and their affective state of trust. By shifting from system-driven personalisation to user-driven moral self-determination, the platform can restore the human element in digital booking, ultimately raising traveller satisfaction (Park & Son, 2025).

Research suggests that good information ethics, combined with a high-quality information system, is a key condition for highly information-literate travellers to achieve high satisfaction. By providing tools that enable filtering based on individual ethical criteria, platforms empower users to make choices independently, thereby fostering deep satisfaction (R. Wang et al., 2024). Conversely, some studies have shown that the lack of transparency of booking platforms regarding their data management or their algorithmic logic can decrease user trust and therefore lead to slower technology adoption. Interestingly, users also tend to respond less negatively to an error if the error was caused by AI instead of humans, suggesting a stronger psychological acceptance of AI-generated recommendations (K. Wang et al., 2023). Ultimately, this multifaceted satisfaction serves as a cornerstone of platform success, directly influencing user loyalty and future behavioural intentions (Ranjbaran et al., 2024).

2.5 Research model and hypotheses development

2.5.1 Conceptual research model

The conceptual framework of this study is anchored in the S-O-R paradigm. Across modern digital environments, researchers have relied on this framework to understand how technology-related inputs impact users' mental and emotional processes, ultimately affecting their behaviour (Park & Son, 2025). Applied to AI-powered travel platforms, it represents the extent to which users can proactively configure the moral, transparency, and data-related parameters that regulate AI-generated travel suggestions. Rather than building on conventional personalisation frameworks, this approach introduces a notably different theoretical standpoint (Zhang & Sundar, 2019). Instead of reiterating the technical or ethical foundations previously detailed in the literature review, this section positions ethical customisation as a human-centred intervention that modifies how travellers interpret and interact with recommendations. The main proposition is that the interaction no longer reflects a passive consumption of automated outputs but becomes a collaborative decision-making process when users gain explicit control over the ethical direction of the AI. This conceptual shift aligns with emerging literature on participatory AI, which emphasises the importance of user empowerment and value alignment (Nizamani et al., 2025). Ethical customisation features also operate as perceptual signals that communicate transparency and user-centricity before any active interaction occurs, reinforcing the stimulus function within the S-O-R framework.

This is particularly relevant to the present study's experimental design, which focuses on the perceptual availability of features rather than their active use. The signal itself, rather than the user's action upon it, constitutes the stimulus. The internal organism stage is represented by three psychological stages, referring to perceived agency, platform trust, and perceived ethical alignment. Agency reflects the user's subjective sense of authorship in the recommendation process; trust reflects the belief that the system behaves transparently and reliably; and perceived ethical alignment reflects the degree to which the user experiences the system's output as morally consistent with their own values. These three internal states jointly determine the response, traveller satisfaction, through distinct but related psychological pathways.

2.5.2 Hypotheses development

2.5.2.1 *Ethical customisation*

The transition from purely automated "Black Box" algorithms to user-governed systems represents a critical evolution in digital travel service design. Recent research suggests that the presence of customisation controls acts as a "psychological buffer" against the perceived invasiveness of AI (Aguirre et al., 2015). This effect stems from the VSD, which argues that aligning technology with individual human morality directly correlates with higher evaluative responses (Del Valle & Lara, 2024). Empirical findings from 2024 indicate that travellers are no longer only satisfied by the functional accuracy of a recommendation, but they increasingly seek moral conformity. This refers to a state where the AI's output confirms the user's ethical identity. When a booking platform provides explicit tools to customise ethical parameters, such as prioritising eco-certified accommodations, it mitigates the choice overload often associated with opaque systems. By allowing for this proactive calibration, the platform shifts from nudging users toward its own benefits to a tool for responsible consumerism. In addition, the implementation of explanatory customisation allows users to see the immediate impact of their ethical choices on the results, creating a feedback loop that strengthens the platform's perceived value. The HITL dynamic ensures that machine intelligence remains an extension of human judgment, leading to a more positive overall evaluation of the digital service encounter (Y. Li et al., 2024). While a direct effect is theoretically plausible, the S-O-R framework primarily anticipates that this relationship will operate through internal organism states.

H1. The perceived availability of ethical customisation features in an AI-powered RS positively and directly influences traveller satisfaction.

2.5.2.2 Mediator: Perceived agency

Contemporary algorithmic systems often suffer from opaque decision-making and automation bias, which together generate a lack of users' perceived transparency and understanding (Parasuraman & Manzey, 2010). Ethical customisation is expected to enhance perceived agency because it returns moral and procedural control to the user. Scholarship in autonomy research demonstrates that feelings of agency are heightened when individuals are granted meaningful input into how systems are designed and operated (Ryan & Deci, 2020). Moreover, emerging literature on algorithmic reflectivity suggests that systems enabling users to adjust the system's internal decision rules reinforce their sense of authorship and significantly strengthen perceived agency.

When users experience system-driven recommendations that repeatedly do not fully align with their identity or values, algorithmic fatigue manifests in lowered engagement, frustration, and a reduced sense of personal control (Yang et al., 2024). Consequently, perceived agency is expected to positively influence traveller satisfaction because agency strengthens the user's sense of involvement and reduces uncertainty in digital decisions.

H2. The relationship between the perceived availability of ethical customisation features and traveller satisfaction is mediated by perceived agency, such that higher perceived agency increases satisfaction.

2.5.2.3 Mediator: Platform trust

AI-powered RSs have been criticised for the potential to engage in nudging practices and to optimise outcomes primarily for platform benefits rather than user interests (Luo et al., 2024). While trust becomes a vulnerable component of the user experience (Jha et al., 2023), ethical customisation makes the platform's intentions visible. Literature on XAI shows that transparency and user-adjustable ethical constraints significantly increase trust, thereby enabling AI to transform into a trusted collaborator (Santosh & Wall, 2022). Trust and fairness perceptions have repeatedly been shown to be mediating mechanisms in digital service evaluations. When users perceive the system as fair, respectful of their values, and trustworthy, their satisfaction increases significantly (Angerschmid et al., 2022). Ethical customisation can therefore drive satisfaction both directly and indirectly, with trust acting as an underlying psychological pathway.

H3. The relationship between the perceived availability of ethical customisation features and traveller satisfaction is mediated by platform trust, such that higher platform trust increases satisfaction.

2.5.2.4 Mediator: Perceived ethical alignment

Beyond agency and trust, the perceived availability of ethical customisation features is expected to act as a third psychological pathway. This concept refers to the sense of moral resonance between the user's values and the system's outputs. When users see that a booking platform offers visible controls related to eco-friendliness, data privacy, or social impact, they are likely to assume that the system's recommendations are shaped by criteria that reflect their own priorities, even before engaging with these features. This assumption forms the basis of this mediator, perceived ethical alignment, which is defined as the subjective sense that an AI system operates in a way that is consistent with users' values (I. Koo et al., 2025). From a VSD perspective, technologies that appear to recognise and incorporate users' moral concerns tend to generate more positive evaluations (Del Valle & Lara, 2024). In AI-powered travel contexts, where concerns about environmental impact, data protection, and social responsibility are increasingly prominent (I. Koo et al., 2025), the visibility of ethical customisation features may act as a signal that user values are taken seriously within the

recommendation process. This sense of being morally recognised or acknowledged by the system is expected to enhance overall satisfaction.

H4. The relationship between the perceived availability of ethical customisation features and traveller satisfaction is mediated by perceived ethical alignment, such that higher value coherence increases satisfaction.

2.5.2.5 Moderator: Algorithmic literacy

While ethical customisation provides the structural opportunity for user governance, its efficiency is likely dependent upon the user's cognitive abilities to navigate complex AI logic. Recent research suggests that without a baseline understanding of how data inputs translate into automated outputs, customisation features may be perceived as a cognitive burden rather than an empowering tool (Oeldorf-Hirsch & Neubaum, 2023). Algorithmic literacy may moderate the relationship between ethical customisation and traveller satisfaction, suggesting that users with a stronger understanding of AI processes can better appreciate the functional and ethical implications of customisation. According to the Cognitive Load Theory, users with a higher algorithmic literacy possess the mental logic necessary to process ethical configuration options without experiencing digital overwhelm (Twabu, 2025). Such users are more capable of interpreting ethical options and assessing their relevance, which can amplify both the perception of fairness and satisfaction derived from the customisation (Xu et al., 2025). Conversely, users with low algorithmic literacy may experience frustration when faced with numerous ethical customisation options (Oeldorf-Hirsch & Neubaum, 2023), since excessive choices can create cognitive overload that reduces satisfaction and diminishes the positive effects of customisation (Long et al., 2025). Literacy also allows users to distinguish between genuine ethical governance and "ethics washing", referring to superficial customisation features designed to mask the actual data collection processes (Cox, 2024).

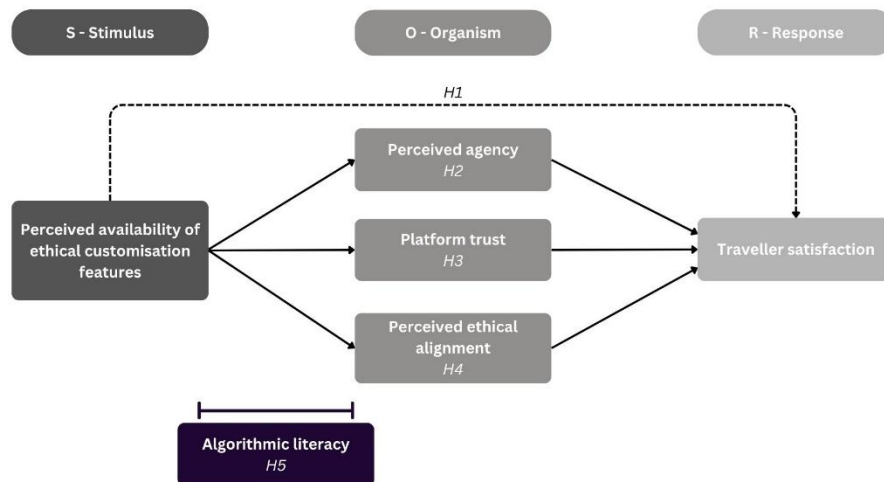
H5. Algorithmic literacy moderates the relationship between the perceived availability of ethical customisation features and traveller satisfaction, such that users with higher algorithmic literacy derive greater satisfaction from ethical customisation features.

2.5.2.6 Control variable: perceived price

To isolate the psychological impact of ethical customisation on traveller satisfaction, this study includes perceived price as a main control variable. Within the tourism industry, price remains the most dominant functional attribute influencing consumer choice and post-purchase evaluation (Riquelme et al., 2021). Because ethical travel options, such as carbon-offsetting or fair-trade accommodations, are frequently associated with a "green premium" or higher operational costs, failure to control for price could result in a biased estimation of the relationship between ethical customisation and satisfaction (Hultman et al., 2015). Economically, price serves as an indicator for the sacrifice a traveller makes in exchange for a service (Dang et al., 2023). According to the Equity Theory, satisfaction is derived from the perceived ratio of outcomes to inputs (K. Kim & Ha, 2024). If a user perceives the price as too high, even a highly ethical AI recommendation may fail to produce satisfaction (Wan & Bharti, 2025). Recent studies in algorithmic pricing suggest that users are increasingly sensitive to "dynamic pricing" models (Alderighi et al., 2022). In the proposed model, perceived price is treated as a control variable to ensure that the observed variance in satisfaction is truly a result of the user's sense of agency and trust rather than a simple reaction to the affordability of the travel suggestions. By eliminating price as a confounding factor, the study has the ability to assess the travellers' enduring preference for moral consistency and control in AI interactions, which extends beyond cost-based motivations (Wan & Bharti, 2025).

2.5.3 Conceptual model

Figure 1
Conceptual Research Model



Source: Author's own elaboration.

The model integrates the S-O-R framework with Signalling Theory to explain how the perceived availability of ethical customisation features influences psychological responses and traveller satisfaction. Arrows represent hypothesised relationships. The dashed arrow indicates the direct path (H1); solid arrows indicate mediated paths (H2-H4); the moderating path (H5) is shown as a bracket under the stimulus-organism segment.

Taken together, the conceptual model positions ethical customisation not as a simple interface feature but as a perceptual intervention that operates simultaneously at three psychological levels. At the level of control, it activates perceived agency by making visible the user's capacity to influence algorithmic logic. At the level of confidence, it builds platform trust by signalling transparent and value-aligned system design. At the level of coherence, it generates perceived ethical alignment by communicating that the platform's recommendation criteria are consistent with the user's moral identity. The S-O-R framework provides the causal architecture connecting these three organism-level states to the behavioural response of traveller satisfaction, while Signalling Theory explains why the visibility of ethical customisation features (independent of their active use) is sufficient to initiate these psychological processes.

Chapter 3: Research design

Through the adoption of a quantitative research design by utilising a between-subjects experiment, an independent variable is deliberately manipulated while participants are distributed across conditions through random assignment (Ramani & Aguinis, 2023). This random distribution between the treatment and control groups removes the influence of third variables, meaning that differences detected in satisfaction can be traced back to the experimental stimulus itself rather than individual differences that already existed (Nguyen, 2025). In particular, the study employs the Experimental Vignette Methodology (EVM), which increases experimental realism by presenting participants with realistic scenarios to evaluate their intentions and attitudes. By building on this approach, ethical customisation mechanisms are manipulated within a tightly controlled experimental setting, which secures internal validity (Aguinis & Bradley, 2014).

3.1 Methodology

The methodology is theoretically grounded in the S-O-R framework. The between-subjects design ensures that this perceptual variable is manipulated through exposure to different interface versions rather than through any active user behaviour. The internal psychological mediation, which acts as the organism in the S-O-R framework, is captured through three mediators, namely the Perceived Agency (the subjective sense of moral self-determination), Platform Trust (the perceived honesty and transparency of the system), and Perceived Ethical Alignment (the degree to which the user experiences the system's outputs as morally consistent with their own values) (I. Koo et al., 2025; Pavlou, 2003). Unlike the experimental manipulation, which is based on how participants are assigned to different interface conditions, the moderation suggested in H5 is tested using an analytical approach. For all participants, algorithmic literacy is treated as a continuous variable, no matter which condition they are in. To assess the moderation effect, interaction terms in regression models are used, which are a standard method for testing moderated mediation in this type of study (Igartua & Hayes, 2021). This method helps maintain statistical power across the entire sample and avoids the issues that come with dividing participants into "high" and "low" literacy categories.

In the final response stage of this framework, traveller satisfaction is used as the key outcome measure, indicating users' concluding emotional reactions to the AI-powered recommendation (Park & Son, 2025). Through a deductive lens, the thesis examines the HITL design philosophy, which positions AI as a complement to human ethical decision-making rather than a substitute (Santosh & Wall, 2022).

3.2 Sample

Prior to data collection, GPower (version 3.1.9.7) was used to run an a priori power analysis, which established the minimum number of participants needed to reliably identify effects that are meaningful. For a two-group between-subjects design using an independent samples t-test, a medium effect size of Cohen's $d=0.50$ was specified as the detection threshold (Sullivan & Feinn, 2012), consistent with benchmarks recommended for experimental vignette studies of this type (Aguinis & Bradley, 2014). With a two-tailed α of .05 and a target power of $1-\beta=.80$, the analysis indicated a minimum of 51 participants per group ($N=102$ total) would be required. The final valid sample exceeded this minimum threshold after the exclusion of failed manipulation checks (which was later revised), ensuring sufficient statistical power for mediation and moderation analyses.

Primary recruitment was conducted via convenience sampling, reaching individuals who are easily accessible through social media and academic networks. To enhance the sample's reach and ensure the inclusion of experienced travellers familiar with algorithm-driven services, the study also incorporated snowball sampling. This technique used referrals from initial subjects to identify hard-to-reach populations, such as frequent travellers who actively seek ethical alignment in their bookings (Nguyen, 2025). All participants provided informed consent before participation, and no personal

identifying information was collected. Additionally, the study was conducted in accordance with the ethical guidelines of HEC Liège and the University of Liège.

Data quality was maintained through an Instructional Manipulation Check. To ensure the quality of data, a symmetric manipulation check was applied across both groups. It should be noted that the application of this exclusion criterion was subsequently revised during the analytical phase. In the treatment group (Group B), participants who answered “No” in the statement “In the version I saw, it appeared that users could adjust how the recommendations were generated” were excluded, as they did not perceive the key experimental feature as intended. In the control group (Group A), participants who answered “Yes” to the equivalent item were also excluded, as they incorrectly perceived customisation features that were not present in their version of the interface. This symmetrical approach ensured that only participants who accurately recognised their assigned condition were included in the final analysis. Note, however, that this planned exclusion strategy was revised during the analytical phase following inspection of response distributions. As detailed in Section 4.2, the manipulation check item appeared to capture not only the recognition of the experimental stimulus but also broader interpretations of algorithmic personalisation, meaning that binary exclusion would have systematically removed participants with higher algorithmic awareness and introduced selection bias correlated with the primary moderator of interest. Retaining the item as a covariate rather than an exclusion criterion was therefore chosen as the methodologically more appropriate choice for this sample and design. This decision is consistent with methodological guidance from Rubio Juan & Revilla (2021), who note that failed manipulation checks reflect meaningful individual differences rather than pure data quality failures when the check item has multiple plausible interpretations. The final wording of the revised item, along with all other measurement instruments, is presented in Section 3.3.

3.3 Research instrument

A structured online questionnaire was developed following Brace (2005), administered in four languages (French, Dutch, German, English) and reviewed for semantic equivalence across translations. The core of the experimental design employs a scenario-based introduction model, which combines controlled laboratory environments with real-world applicability as argued by Jasso (2006). In this study, participants were prompted to imagine a specific travel context that established a uniform frame of reference across the sample. Booking.com was selected as the platform stimulus because it is one of the most widely used online accommodation booking platforms in Europe and worldwide, ensuring high ecological validity and recognition among the target sample drawn from various countries.

The decision to use a static interface image rather than a live interactive prototype is a deliberate and theoretically grounded methodological choice. Prior research on initial trust formation shows that users form stable judgements about digital platforms within seconds of first visual exposure, often even before engaging with any features (Pavlou, 2003). In this study, the focus is specifically on the perceptual effects of ethical customisation availability, rather than the behavioural consequences of actively using those features. For this reason, a static image provides the most appropriate experimental stimulus. This approach is also consistent with the EVM, which supports the use of realistic but controlled scenarios to obtain meaningful attitudinal and evaluative responses (Aguinis & Bradley, 2014). By presenting participants with a carefully designed booking interface image, the study isolated the effect of perceived feature availability while holding all other elements constant. This way, the study strengthens internal validity and keeps the causal interpretability tightly focused on perception rather than behaviour.

Following this introduction, a randomised stimulus-response model was implemented (Ihrke, 2011). Randomisation serves as a fundamental procedural safety measure, ensuring that any observed variances in outcomes are a direct result of the experimental manipulation rather than confounding variables (Onghena, 2017). At the beginning of the survey, all participants were asked to imagine they

were logged into Booking.com and were planning a three-day trip to Brussels, Belgium (search parameters: Friday 16 October- Sunday 18 October, 1 adult, 1 room). They were told that the platform's AI recommendation system had generated hotel suggestions based on their profile, previous activity, and general travel trends. Participants were then randomly assigned to one of two conditions. Group A (control group) viewed a standard Booking.com hotel search interface, displaying three hotel results: Hotel City Central (145€), Urban Stay Hotel (148€), and Metro Plaza Hotel (146€). The only functional attributes that were visible are the price, the ratings, and the number of reviews. Group B (treatment group) viewed the same interface with the addition of three ethical sliders: Eco-Friendly (with the description: level of priority given to hotels with recognised eco-labels and active waste-reduction programs), Data Privacy (with the description: personalisation is based only on the current session; identity is not shared or tracked externally), and Social Impact (with the description: boosts hotels that source at least 50% of their labour and supplies from local businesses). All sliders were pre-set to Medium, with explanatory tooltips visible.

The analytical framework of this thesis relied on closed-ended, fixed-response questions. This structured approach is essential for this quantitative survey, since it provides a standardised environment for comparing participant perceptions and performing accurate statistical analysis (Malhotra et al., 2017). All psychological constructs were measured using validated 7-point Likert scales, ranging from (1) "Strongly Disagree" to (7) "Strongly Agree". The measurement scales were adapted from established literature to ensure construct validity.

- **Perceived Agency:** This construct was measured using three items that were adapted from Jannach et al. (2017), capturing users' perceived ability to influence recommendation generation, exercise control over outputs, and effectively express their preferences (see Appendix A for full item wording). Items were selected based on their alignment with autonomy research demonstrating that meaningful input into system design heightens perceived agency (Ryan & Deci, 2020).

Considering the AI-powered recommendation system on this booking platform:

- "It appears to give users the ability to influence how the AI-powered recommendations are generated."
- "It appears to allow users to be in control of the recommendations."
- "It would allow me to effectively express my preferences."

- **Perceived ethical alignment:** This construct was measured using three items adapted from I. Koo et al. (2025) by focusing on value coherence between the user's moral priorities and the system's visible outputs (see Appendix A). These items were selected because they capture the personalised experience of moral resonance central to this study's operationalisation of the construct, rather than abstract beliefs about whether AI systems follow ethical principles in general.

Considering the recommendations generated by the AI-powered recommendation system on this booking platform:

- "They seem to match my personal values."
- "They seem to reflect what matters to me when choosing travel options."
- "They seem to take into account aspects that are important to me personally."

Although these items were initially created to measure ethical AI adoption, they were chosen for this study because they focus on the key aspect of value coherence, meaning how much the system's results align with the user's moral values. This alignment is central to how ethical alignment is perceived in this thesis.

- **Platform trust:** This concept was measured using three items capturing competence, benevolence, and transparency dimensions of trust, respectively (Afroogh et al., 2024; Jiang et al., 2025) (see Appendix A). This is consistent with the multidimensional trust frameworks applied in AI-mediated service contexts.

Considering the AI-powered recommendation system on this platform:

- “I trust it to provide reliable recommendations.”
 - “I believe it acts in my best interest.”
 - “It operates in a transparent way.”
- **Traveller satisfaction:** Three items were adapted from Yu et al. (2024) by relying on the Expectation Confirmation Model (ECM) and Technology Acceptance Model (TAM) (see Appendix A). This thesis suggests that AI user retention is driven by the confirmation of expectations, which turns into satisfaction, ultimately being a critical predictor of long-term platform loyalty.
 - “I would be satisfied with the AI-powered recommendations on this booking platform.”
 - “The AI-powered recommendations on this booking platform seem to meet my expectations.”
 - “Using this booking platform, supported by its AI-powered recommendation system, provides a pleasant user experience.”

A fourth item functions as a secondary, exploratory outcome item measuring behavioural intention, consistent with the Expectation Confirmation Model (ECM) and Technology Acceptance Model (TAM) logic (Yu et al., 2024). Meanwhile, satisfaction and behavioural intention are treated as conceptually distinct. Satisfaction is the primary dependent variable of this thesis, while behavioural intention is included as an unlabelled downstream indicator to enable explanatory analysis. It is not formally hypothesised as a separate outcome.

- “I could imagine myself using a booking platform that includes an AI-powered recommendation system like this in a real booking situation.”
- **Algorithmic literacy:** The scale items for algorithmic literacy were adapted from Dogruel et al. (2022), operationalised as perceived rather than objective knowledge to capture booking-platform-specific understanding of how AI recommendation systems personalise suggestions and use personal data (see Appendix A). Perceived understanding was prioritised over objective knowledge testing because it more directly predicts satisfaction and behavioural responses in AI-mediated contexts (Wu & Xie, 2024).
 - “I am familiar with how AI-powered recommendation systems personalise suggestions on booking platforms.”
 - “I feel confident in my understanding of how AI-powered recommendation systems work on booking platforms.”
 - “I have a general understanding of how my data is used by AI-powered recommendation systems on booking platforms.”
 - **Price-ethical trade-off orientation:** These two conceptually distinct single items measured participants’ ethical considerations over price, as well as their willingness to pay a premium. They were included separately as covariates to control for individual differences, ensuring that any observed variation in satisfaction could be attributed to the ethical customisation

manipulation rather than respondents' pre-existing attitudes toward price or ethical consumption (Hultman et al., 2015; Sun & Yoon, 2022).

- "I prioritise ethical aspects over price when booking."
- "I am willing to pay a higher price for a hotel that meets my personal expectations."
- **Ethical orientation (pre-exposure baseline):** Three items were chosen to measure the personal importance of eco-friendliness, community support, and responsible data handling when making travel decisions. These items were rated on the same 7-point Likert scale (1= Strongly Disagree, 7= Strongly Agree) to capture these values, which align with the themes explored later in the experiment (M. Koo & Yang, 2025). First, they help ensure that the two experimental groups were similar in terms of their ethical orientation before being exposed to the manipulation. Second, they are used as potential control variables in the main analysis to account for any pre-existing ethical attitudes, so the results of the manipulation on satisfaction and other factors are not influenced by these individual differences. While these items are not included in the hypothesised model, they serve as a pre-exposure control measure and an equivalence check.
 - "It is important to me that travel options are environmentally friendly."
 - "It is important to me that travel options support local communities."
 - "It is important to me that my personal data is handled responsibly."

Demographic data (age, gender, country of origin, and current employment status) were collected at the end of the survey. These variables are used to provide a basic overview of sample composition and to support simple statistical controls where necessary. Travel frequency was captured earlier in the surveys (within the experience section) rather than in the demographic block, since it was treated as a pre-exposure experiential factor that may influence how participants interpret the interface, rather than as a pure background characteristic.

3.4 Pilot test

Before distributing the final survey instrument, an initial rollout of the questionnaire revealed substantial participant confusion and critical feedback regarding the clarity of the instructions, the framing of the experimental vignette, and the interpretation of several scale items. As a result, data collection was temporarily halted, and the instrument underwent a structured revision phase with a convenience sampling of 24 respondents. Wording adjustments were made to improve the clarity of the ethical slider descriptions, and several scale items were simplified to reduce ambiguity across translated versions. Feedback from pilot participants also led to small layout modifications intended to improve readability and survey flow.

The most consequential revision concerned the manipulation check item, which was redesigned from a four-option recall test to a binary Yes/No awareness measure administered to both groups. In the final instrument, the item was replaced with a Yes/No awareness measure administered to both groups: "In the version I saw, it appeared that users could adjust how the recommendations were generated."

The stimulus framing was also substantially revised. In the initial version, treatment-group participants were explicitly instructed to "evaluate the specific features that allow you to personally adjust the AI's logic using ethical sliders" and to "observe how these recommendations align with your personal values." Feedback from the halted rollout and subsequent pilot indicated that these instructions felt overly prescriptive, and preliminary item analysis suggested that they may have artificially inflated perceived agency scores by directing participants' attention toward the intended construct. To address this issue, the final stimulus adopted a more neutral framing across both conditions.

Participants in both groups received the same general observation instructions, while the treatment condition received only one additional sentence informing participants that certain parameters could be adjusted. This change reduced experimental demand bias and more efficiently isolated the perceptual effect of ethical customisation rather than participants' deliberate evaluation of it.

Several revisions were also made at the scale level following item analysis and content validity assessment. The perceived agency scale was reduced from four items to three by removing the past-tense ownership statement, "The final choices felt like they were mine", because it conceptually overlapped with the satisfaction construct and risked cross-construct contamination. The eco-ethical value items were rephrased from behavioural-frequency wording (e.g., "I actively seek out eco-friendly options") to importance-based wording (e.g., "It is important to me that travel options are environmentally friendly"). This adjustment reduced the fusion of value orientation with past behaviour and lowered sensitivity to social desirability bias. In addition, the algorithmic literacy scale was expanded from two items to three and shifted from measuring general familiarity with AI tools to capturing booking-platform-specific cognitive understanding. This revision better reflected the construct's relevance to the study's moderation hypothesis.

Two broader instrument-level changes were implemented as well. First, an instructed-response item was embedded within the platform trust scale, asking participants to select "Agree". This item was inadvertently omitted from the final questionnaire during the revision process. To maintain data quality despite its removal, response screening was instead conducted through objective time-based exclusion criteria applied during the analysis phase (see Section 4.1). Second, the hotel stimuli were revised to minimise price-related confounding effects. The original interface displayed hotel prices ranging from €80 to €110, a variation large enough to independently influence hotel preference and satisfaction. In the final version, prices were standardised to a near-identical range (€145, €146, €148), thereby neutralising price as a competing explanatory factor and ensuring that any observed differences in satisfaction could be more confidently attributed to the ethical customisation manipulation itself.

Chapter 4: Results

The following chapter reports the empirical results of the between-subjects experiment investigating whether users' perceived access to ethical customisation features in an AI-powered travel booking platform influences traveller satisfaction, and through which psychological mechanisms. Results are reported in six sequential sections: (1) preliminary scale reliability and structural validity analyses, (2) sample description and randomisation checks, (3) descriptive statistics and correlations, (4) hypothesis testing, (5) exploratory analysis for the scale item Behavioural Intention, and (6) supplementary insights from qualitative comments.

4.1 Preliminary Analysis: Scale Reliability and Structural Validity

4.1.1 Internal Consistency Reliability

Before testing the hypotheses, the internal consistency of all multi-item scales was evaluated using Cronbach's alpha (α). Omega calculations via CFA were not estimable for all scales due to sample size constraints, and therefore, alpha is reported throughout. As a general benchmark, $\alpha \geq .70$ is considered acceptable, $\alpha \geq .80$ is considered good, and $\alpha \geq .90$ is considered excellent (Nunnally, 1978). Table 1 summarises the reliability estimates for each scale.

Table 1

Internal Consistency Estimates for All Study Scales

Scale	Items	Cronbach's α	95% CI Lower	95% CI Upper
Price-ethical Trade-off (PS)	PS1_r, PS2_r	0.336*	—	—
Perceived Agency (PA)	PA1_r, PA2_r, PA3_r	0.804	0.728	0.881
Ethical Alignment (EA)	EA1_r, EA2_r, EA3_r	0.915	0.863	0.967
Platform Trust (PT)	PT1_r, PT2_r, PT3_r	0.783	0.653	0.912
Satisfaction (S)	S1_r, S2_r, S3_r	0.894	0.855	0.933
Algorithmic Literacy (AL)	AL1_r, AL2_r, AL3_r	0.801	0.730	0.871
Eco-ethical Values (EV)	EF1_r, EF2_r, EF3_r	0.743	0.658	0.828

*Note. * Price-ethical trade-off orientation is a two-item scale (PS1_r and PS2_r). Due to the minimum of three items required for if-item-dropped statistics and omega estimation, only alpha is reported. The low α reflects the items' conceptual distinctiveness; they are therefore entered individually as covariates in all models rather than as a composite. CI = confidence interval. Full item-dropped statistics and omega output for all scales are presented in Appendix C1.*

Source: Author's own elaboration based on JASP output.

All scales achieved acceptable or good internal reliability. The Ethical Alignment ($\alpha = 0.915$), the Satisfaction ($\alpha = 0.894$), the Perceived Agency ($\alpha = 0.804$), and the Algorithmic Literacy ($\alpha = 0.801$) scales all demonstrated excellent consistency. The Platform Trust ($\alpha = 0.783$) and Eco-ethical Values ($\alpha = 0.743$) scales all exceeded the .70 threshold comfortably. The price-ethical trade-off orientation scale, consisting of only two items, returned $\alpha = 0.336$, which falls below acceptable benchmarks. However, this two-item scale is included solely as a control variable. Given the very low inter-item correlation, PS1_r (ethics over price) and PS2_r (willingness to pay a premium) are better understood as conceptually distinct single-item indicators than as a composite scale and are

therefore entered individually in all models accordingly. Overall, the scales demonstrated acceptable internal consistency for subsequent analyses.

4.1.2 Structural validity: Exploratory Factor Analysis

To assess the structural validity of the adapted scales, an exploratory factor analysis (EFA) was conducted on the 15 items representing the five main constructs: Perceived Agency (PA), Ethical Alignment (EA), Platform Trust (PT), Satisfaction (S), and Algorithmic Literacy (AL). Control variables were excluded. Before extracting factors, the data's suitability was confirmed: Bartlett's test of sphericity was significant ($\chi^2 (105) = 1188.046, p < .001$), and the Kaiser-Meyer-Olkin (KMO) measure was 0.868, indicating that the correlation matrix was well-suited for factor extraction. Instead of the planned five categories, the results revealed a three-factor solution that explained 60.9% of the total variance. While Perceived Agency and Algorithmic Literacy emerged as distinct and clear factors, the items for Ethical Alignment, Platform Trust, and Satisfaction all merged into a single factor.

Specifically, Factor 1 comprised all items from Ethical Alignment, Platform Trust, and Satisfaction; Factor 2 comprised the three Perceived Agency items; and Factor 3 comprised the three Algorithmic Literacy items. This happened because those three constructs were very highly correlated (r between 0.652 and 0.721), which substantially exceeded their correlations with Perceived Agency ($r = 0.425$ - 0.558). In a single-exposure experiment, where participants evaluate a static interface image, responses to evaluative constructs such as trust, value alignment, and satisfaction are likely driven by a shared global evaluation of the stimulus, rather than by distinct psychological processes, suggesting that full discriminant validity between these three constructs cannot be assumed in a single-exposure experimental context (which is discussed further in Section 5.6).

Table 2

EFA: Factor Loadings for Main Study Constructs (N = 122)

Item	Factor 1: Evaluative Response	Factor 2: Perceived Agency	Factor 3: Algorithmic Literacy
EA1_r	0.677		
EA2_r	0.661		
EA3_r	0.651		
PT1_r	0.690		
PT2_r	0.869		
PT3_r	0.589		
S1_r	0.851		
S2_r	0.846		
S3_r	0.700		
PA1_r		0.707	
PA2_r		0.820	
PA3_r		0.632	
AL1_r			0.681

AL2_r			0.679
AL3_r			0.885
Eigenvalue	6.550	1.536	1.045
% Variance	0.437	0.102	0.070
Cumulative %	0.437	0.539	0.609
<i>Note. N = 122. Principal axis factoring with oblique rotation. Loadings below .40 suppressed for clarity. Factor 1 comprises Ethical Alignment, Platform Trust, and Satisfaction items, consistent with the high inter-construct correlations observed ($r = 0.65\text{--}0.72$). Bartlett's test of sphericity: $\chi^2(105) = 1188.046$, $p < .001$; KMO = 0.868.</i>			

4.1.3 Supplementary Confirmatory Evidence: Single-factor CFA

To further validate the scales, separate Confirmatory Factor Analysis (CFA) models were run for each construct using the Weighted Least Squares Mean and Variance adjusted estimator (WLSMV), which is appropriate for ordinal Likert-scale data (C.-H. Li, 2016). This approach was adopted in place of a full simultaneous measurement model, which could not be stably estimated given the sample size relative to the number of free parameters required. Because these models used only three items each, they represent a saturated model, which is why perfect fit indices are expected. The substantive evidence, therefore, lies in the significance and direction of the factor loadings and the admissibility of the solution (positive residual variances). Results for Perceived Agency, Ethical Alignment, Satisfaction, and Algorithmic Literacy all produced admissible solutions with all factor loadings statistically significant (all $p < .001$) and all residual variances positive, which confirms that each set of items measures a single underlying construct.

While the Platform Trust scale demonstrated acceptable internal consistency ($\alpha = 0.783$, 95% CI [0.653, 0.912]), the model produced an inadmissible solution due to a negative residual variance estimate (a Heywood case). A Heywood case arises when one item's variance is almost entirely explained by the common factor, leaving a residual that the estimator rounds to zero or below. This is consistent with the item-level statistics for item PT3_r. The item's α -if-dropped of 0.775 is only marginally below the full-scale α of 0.783, indicating that its variance is almost fully accounted for by the shared factor (see Appendix C1.4). By contrast, PT2_r has an α -if-dropped of 0.508, confirming that it carries substantial shared variance with the other items and that the construct is internally coherent. PT1_r has an α -if-dropped of 0.798, marginally above the full scale, suggesting that it is the least redundant of the three items. These patterns indicate that the scale items measure a common dimension but with unequal communality profiles that a saturated three-item structure cannot simultaneously accommodate. The Heywood case is therefore attributed to model identification constraints rather than construct invalidity. The composite scale was retained for all subsequent analyses, with this limitation noted.

Table 3

Single-Factor CFA Results: Standardised Factor Loadings by Construct (N = 122)

Construct	Item	Standardised Loading	SE	p
Ethical Alignment	EA1_r	1.000	0.000	/
	EA2_r	1.083	0.052	<.001
	EA3_r	0.942	0.030	<.001

Platform Trust	PT1_r	The model is not admissible: lavaan- >lav_object_post_check(): some estimated ov variances are negative		
	PT2_r			
	PT3_r			
Satisfaction	S1_r	1.000	0.000	/
	S2_r	1.052	0.062	<.001
	S3_r	0.836	0.039	<.001
Perceived Agency	PA1_r	1.000	0.000	/
	PA2_r	1.250	0.094	<.001
	PA3_r	1.028	0.067	<.001
Algorithmic Literacy	AL1_r	1.000	0.000	/
	AL2_r	1.079	0.083	<.001
	AL3_r	1.249	0.107	<.001

Note. N = 122. Separate single-factor CFA models estimated per construct using the WLSMV estimator (ordinal data). First item of each construct fixed to 1.0 for identification (marker variable method); SE and p are therefore not reported for reference items. All reported loadings were significant at $p < .001$. The Platform Trust model produced a Heywood case (negative residual variance on PT3_r); this solution is reported for completeness but interpreted with caution.

4.2 Sample description and Randomisation Check

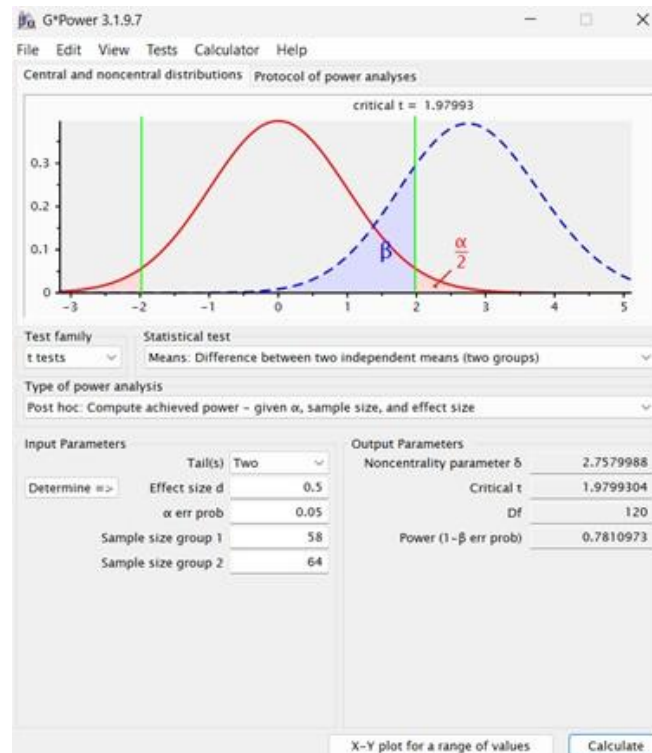
A total of 223 individuals initiated the online survey, of whom 131 reached the final page and were classified as completers. In addition, two respondents who did not formally complete the survey were retained in the dataset because they progressed sufficiently far into the questionnaire to provide usable responses: one reached page 9 (did not complete the price sensitivity items) and the other page 11 (did not complete/ passed the optional free comment section at the end) of the survey flow.

Three data quality filters were then applied sequentially to the 133 usable responses (131 completers plus two partial completers who provided sufficient data). First, respondents whose total completion time exceeded 60 minutes ($n=4$) were excluded, as an extended absence between pages can compromise controlled exposure to the experimental stimulus. Second, respondents whose total completion time fell below one-third of the median completion time (6 minutes 36 seconds; 132 seconds) were excluded ($n=1$) to remove potentially insufficient effort responses (Meade and Craig, 2012). Third, respondents who spent fewer than five seconds on the scenario page containing the experimental stimulus image ($n=5$) were excluded, as this duration is insufficient for meaningful visual processing of a complex booking interface. One additional participant in the treatment condition was excluded following an explicit post-completion withdrawal of consent. These sequential exclusions yielded a final analytical sample of $N=122$ ($133-4-1-5-1=122$), comprising 58 participants in the control condition and 64 in the treatment condition.

A post hoc power analysis was subsequently conducted in G*Power (version 3.1.9.7) to determine the statistical power achieved with this final unequal sample. Using the same parameters as the a priori analysis ($d=0.50$, $\alpha=.05$, two-tailed), the post hoc calculation yielded an achieved power of 0.781 (critical $t=1.980$, $df=120$, non-centrality parameter $\delta=2.758$), as illustrated in Figure 2 (Faul et al., 2007). Although this falls marginally below the conventional 0.80 threshold, it closely approaches it

and is considered acceptable for a study of this scope and design. The minor shortfall is attributable to the slight reduction in sample size following the response time exclusions.

Figure 2



*Post hoc power analysis output (G*Power 3.1.9.7) for the final sample ($N = 122$, $d = 0.50$, $\alpha = .05$, two-tailed). Output generated by the author using G*Power (Faul et al., 2007).*

The sample was drawn primarily from Belgium, Germany, Luxembourg, and surrounding regions. In the control group, 84.5% of participants were Belgian, compared to 82.8% in the treatment group. The remaining participants were distributed across Germany, Luxembourg, and a small number of other countries, reflecting the multilingual distribution strategy of the survey.

Gender distributions were broadly comparable across conditions. The control group consisted of 62.1% female and 37.9% male participants, while the treatment group was 73.4% female, 23.4% male, and 1.6% preferring not to say. Age distributions differed somewhat between conditions. The control group was most heavily represented in the 45-54 age bracket (34.5%), whereas the treatment group had its largest proportion in the 55-64 range (26.6%). Both groups were predominantly composed of full-time and part-time employees, and most of the respondents reported prior experience with AI-powered booking platforms.

To verify that the randomisation created comparable groups in terms of key pre-existing individual differences, an independent-samples t-test was performed on algorithmic literacy scores. The results showed a statistically significant difference between the groups ($t(120) = 2.108$, $p = .037$, $d = 0.382$), with the control group reporting modestly higher algorithmic literacy ($M = 4.293$, $SD = 1.439$) than the treatment group ($M = 3.745$, $SD = 1.431$). Although random assignment was applied, this difference likely reflects natural sampling variability, given the moderate sample size. Three points mitigate the interpretive threat this poses. First, because algorithmic literacy is included as a covariate in all subsequent regression and mediation models, its influence on satisfaction is statistically controlled regardless of its between-group distribution. Second, individuals with higher levels of algorithmic literacy may evaluate algorithmic systems more critically, including scepticism toward surface-level ethical framing. If so, the higher literacy observed in the control group would be unlikely to

artificially inflate positive treatment effects. Third, the moderation analysis explicitly models the interaction between group condition and algorithmic literacy, meaning that the differential distribution is incorporated into rather than ignored by the primary analytical model. The imbalance is nonetheless acknowledged as a limitation in Section 6.4.

Although Section 3.2 initially proposed a symmetric exclusion strategy based on the manipulation check item, this procedure was revised during the analytical stage after inspection of the response distributions and participants' comments. The check appeared to capture not only recognition of the experimental stimulus, but also broader interpretations of algorithmic personalisation and recommendation transparency, meaning that binary exclusion risked systematically removing participants with higher algorithmic awareness. Retaining the item as a covariate, therefore, allows us to preserve the full analytical N, avoid selection bias, and control for manipulation awareness across all subsequent models. This deviation from the pre-planned strategy is acknowledged as a limitation in Section 6.4.

4.3 Descriptive Statistics and Correlations

Mean scores ranged from 3.948 (Algorithmic Literacy) to 4.495 (Ethical Alignment), indicating moderate to moderately high levels across constructs. Distributions were approximately normal, with acceptable skewness and kurtosis values.

Table 4

Descriptive Statistics for Principal Constructs

Variable	N	M	SD	Skewness	Kurtosis	Min	Max
1. Perceived Agency	122	4.391	1.380	-0.520	-0.532	1.00	7.00
2. Ethical Alignment	122	4.495	1.376	-1.073	0.286	1.00	6.00
3. Platform Trust	122	3.948	1.261	-0.421	-0.430	1.00	6.33
4. Satisfaction	122	4.448	1.312	-0.824	-0.059	1.00	6.67
5. Algorithmic Literacy	122	4.005	1.455	-0.157	-0.761	1	7
6. Eco-ethical Values	122	4.995	1.187	-0.752	0.659	1	7

Note. N = 122. All constructs were measured on 7-point Likert scales (1 = Strongly Disagree to 7 = Strongly Agree).

Source: Author's own elaboration based on JASP output.

The correlation matrix revealed Ethical Alignment having the strongest correlation with Satisfaction ($r = 0.721$, $p < .001$), followed by Platform Trust ($r = 0.710$, $p < .001$) and Perceived Agency ($r = 0.463$, $p < .001$). Perceived Agency and Ethical Alignment were substantially intercorrelated ($r = 0.558$, $p < .001$), consistent with the theoretical expectation that these organism-level constructs co-activate in response to ethical interface stimuli. Eco-ethical values orientation showed a significant negative correlation with satisfaction ($r = -0.219$, $p = .015$), a counterintuitive pattern suggesting that the more strongly users care about eco-ethical values, the less satisfied they are with the system, possibly because they are more critical or have higher expectations of ethical performance (interpreted fully in Section 6.2). Algorithmic Literacy correlated positively with Satisfaction ($r = 0.328$, $p < .001$) and with all other key constructs. Platform Trust and Ethical Alignment were also strongly intercorrelated ($r = 0.652$, $p < .001$), the highest inter-construct correlation observed, which suggests partial conceptual overlap between these two organism-level variables (see full Pearson correlation matrix in Appendix C6).

Table 5

Pearson Correlation Matrix with Reliability Coefficients on the Diagonal

Variable	1	2	3	4	5	6	α
1. Perceived Agency	—						0.804
2. Ethical Alignment	0.558***	—					0.915
3. Platform Trust	0.425***	0.652***	—				0.783
4. Satisfaction	0.463***	0.721***	0.710***	—			0.894
5. Algorithmic Literacy	0.190*	0.292**	0.221*	0.328***	—		0.801
6. Eco-ethical Values	-0.063	-0.219*	-0.127	-0.219*	-0.271**	—	0.743

Note. $N = 122$. Cronbach's α on the diagonal. *** $p < .001$, ** $p < .01$, * $p < .05$.

Source: Author's own elaboration based on JASP output.

4.4 Hypothesis Testing

4.4.1 H1: Direct Effect of Ethical Customisation on Satisfaction

H1 proposed that the perceived availability of ethical customisation features would directly and positively predict traveller satisfaction. Satisfaction was operationalised as the mean of three items (S1_r, S2_r, S3_r), while BI1 was retained separately as an exploratory behavioural intention indicator. An independent-samples t-test revealed no significant difference in Satisfaction between the control group ($M=4.529$, $SD=1.414$) and the treatment group ($M=4.375$, $SD=1.219$): $t(120) = 0.645$, $p = .520$, $d = 0.117$.

Table 6

Independent Samples t-Tests: Group Comparisons Across All Constructs

Construct	Control M (SD) n = 58	Treatment M (SD) n = 64	t (df)	p	Levene's p	d	Hypothesis
Perceived Agency	3.989 (1.402)	4.755 (1.264)	-3.176 (120)	.002**	.464	-0.576	H2 Step 1 ✓
Ethical Alignment	4.362 (1.483)	4.615 (1.270)	-1.013 (120)	.313	.056	-0.184	H4 Step 1 X
Platform Trust	4.086 (1.297)	3.823 (1.224)	1.153 (120)	.251	.581	0.209	H3 Step 1 X
Satisfaction	4.529 (1.414)	4.375 (1.219)	0.645 (120)	.520	.218	0.117	H1 X

Algorithmic Literacy	4.293 (1.439)	3.745 (1.431)	2.108 (120)	.037*	.713	0.382	Covariate
-----------------------------	---------------	---------------	-------------	-------	------	-------	-----------

*Note. N = 122. Student's t-test reported for all constructs as Levene's Test of Equality of Variances $p > .05$. Cohen's $d = (M_{\text{Control}} - M_{\text{Treatment}}) / \text{pooled SD}$. A positive d value indicates the control group scored higher, since a negative d value indicates the treatment group scored higher. For Perceived Agency, the negative d confirms that the treatment group reported higher agency scores, consistent with H2. For Platform Trust, the positive d indicates control group trust was marginally higher, contrary to the hypothesised direction for H3. ** $p < .01$, * $p < .05$.*

Source: Author's own elaboration.

A hierarchical linear regression further examined the direct effect while controlling for individual differences. Group membership alone explained less than 1% of the variance in Satisfaction and remained non-significant after inclusion of covariates (Estimate=0.009, $p=.971$; see Table 7). A parallel estimate from PROCESS Model 1, which includes the Group x Algorithmic Literacy interaction term, produced a comparable non-significant coefficient (Estimate=0.007, $p=.977$; see Table 8), confirming the robustness of this null result across model specifications. The full regression model was significant, $F(6,114) = 3.400$, $p=.004$, explaining 15.2% of variance in Satisfaction ($R^2=0.152$) (see Table 7). Algorithmic Literacy emerged as the only significant positive predictor (Estimate=0.262, $\beta=.085$, $p=.003$), while the remaining predictors were non-significant. **H1 was therefore not supported.**

Table 7

Hierarchical Linear Regression Predicting Traveller Satisfaction (S_comp)

Predictor	b	SE	β	t	p	$R^2 / \Delta R^2$
Block 1: Group_d						$R^2 = 0.002$, $p = .887$
Group (1 = Treatment)	-0.116	0.237		-0.489	.625	
Block 2: + Covariates						$R^2 = 0.152$, $p = .004$
Group (1 = Treatment)	0.009	0.236		0.036	.971	
Include (awareness)(1=correct answer)	0.013	0.257		0.050	.960	
EV_comp	-0.107	0.117	-0.096	-0.909	.366	
AL_comp	0.262	0.085	0.294	3.090	.003**	
PS1_r	-0.121	0.083	-0.152	-1.465	0.146	
PS2_r	-0.038	0.100	-0.036	-0.383	.703	

Note. N = 122 (one missing on Include). *B* reported for continuous predictors only. Adjusted $R^2 = 0.132$.
 ** $p < .01$.

Source: Author's own elaboration.

4.4.2 H2: Mediation by Perceived Agency

H2 proposed that the effect of ethical customisation availability on satisfaction would be mediated by Perceived Agency. The first condition for mediation (a significant effect of the manipulation on the mediator) was confirmed. Treatment-group participants reported significantly higher Perceived Agency than control-group participants ($M_{\text{Treatment}}=4.755$, $SD=1.264$ vs $M_{\text{Control}}=3.989$, $SD=1.402$; $t(120) = -3.176$, $p=.002$, $d=-0.576$).

To further examine whether perceived agency explained the relationship between ethical customisation and traveller satisfaction, a mediation analysis was conducted using structural equation modelling. The model assessed the indirect influence of experimental group assignment (Group_d) on satisfaction through perceived agency while controlling for perceived price fairness and eco-ethical values. The results demonstrated a statistically significant indirect effect of Group_d on satisfaction through Perceived Agency (Estimate=0.364, $SE=0.124$, $z=2.932$, $p=.003$, 95% CI [0.135, 0.680]). The confidence interval did not include zero, indicating evidence for a significant indirect mediation effect. This indicates that users' Perceived Agency was considerably raised by exposure to ethical customisation features, which in turn had a positive impact on Satisfaction. The direct effect of group condition on satisfaction remained significant and negative (Estimate=-0.505, $SE=0.212$, $z=-2.383$, $p=.017$, 95% CI [-0.912, -0.088]), whereas the total effect of group condition on satisfaction was not statistically significant (Estimate=-0.141, $SE=0.226$, $z=-0.624$, $p=0.533$, 95% CI [-0.618, 0.307]). This pattern (a significant positive indirect effect coexisting with a significant negative direct effect and a non-significant total effect) is an inconsistent mediation pattern, described as suppressor-type mediation (MacKinnon et al., 2000). The total effect is effectively cancelled by the opposite direct and indirect paths, meaning that ethical customisation's positive influence on satisfaction operates exclusively through the agency pathway, while the direct exposure to ethical controls appears to exert a counteracting negative pressure on satisfaction when agency is held constant. The theoretical interpretation of this suppressor effect is discussed in Section 5.2.

The path coefficients further demonstrated that group condition significantly predicted Perceived Agency (Estimate=0.796, $SE=0.244$, $z=3.263$, $p=.001$, 95% CI [0.316, 1.276]), while Perceived Agency significantly predicted satisfaction (Estimate=0.457, $SE=0.088$, $z=5.172$, $p<.001$, 95% CI [0.273, 0.619]). Among the control variables, Eco-ethical Values negatively predicted satisfaction (Estimate=-0.202, $SE=0.101$, $z=-1.998$, $p=.046$, 95% CI [-0.385, 0.018]), whereas perceived price fairness variables showed no significant effects (see full mediation output in Figure C8.7). **H2 is supported.**

4.4.3 H3: Mediation by Platform Trust

H3 proposed mediation through Platform Trust. Contrary to the hypothesised direction, Platform Trust was directionally lower in the treatment group ($M=3.823$, $SD=1.224$) than in the control group ($M=4.086$, $SD=1.297$), though this difference did not reach statistical significance ($t(120) = 1.153$, $p=.251$, $d=0.209$).

To further test this hypothesis, a mediation model was examined. The indirect effect through Platform Trust was not statistically significant (Estimate= -0.177, $SE=0.158$, $z=-1.120$, $p=.263$, 95% CI [-0.508, 0.145]). The direct effect of group condition on satisfaction was also not significant (Estimate= 0.036, $SE=0.173$, $z=0.205$, $p=.837$, 95% CI [-0.294, 0.397]), as was the path from group condition to Platform Trust (Estimate= -0.249, $SE=0.229$, $z=-1.084$, $p=.278$, 95% CI [-0.691, 0.206]). Although Platform Trust significantly predicted Satisfaction (Estimate=0.710, $SE=0.067$, $z=10.66$, $p<.001$, 95% CI [0.566, 0.828]), ethical customisation did not significantly influence trust levels.

As in the previous model, eco-ethical values demonstrated a marginally significant association with satisfaction (Estimate= -0.177, SE=0.098, $z=-1.805$, $p=.071$, 95% CI [-0.354, 0.027]), while the remaining control variables were not significant (see full mediation output in Figure C8.8). **H3 is therefore not supported.** It is further noted that the EFA results reported in Section 4.1.2 indicated that Platform Trust items loaded into the same factor as Ethical Alignment and Satisfaction, suggesting limited discriminant validity between these constructs in a single-exposure design. This should be considered when interpreting the null results.

4.4.4 H4: Mediation by Ethical Alignment

H4 proposed mediation through Perceived Ethical Alignment. The group difference in Ethical Alignment was in the hypothesised direction ($M_{\text{Treatment}}= 4.615$, $SD= 1.270$ vs $M_{\text{Control}}=4.362$, $SD=1.483$) but did not reach statistical significance ($t(120) = -1.013$, $p=.313$, $d=-0.184$), failing the first condition for mediation. Although Ethical Alignment is the strongest within-sample predictor of Satisfaction ($r=0.721$, $p<.001$, see Table 5), the single-exposure manipulation was insufficient to produce a significant between-group difference in this construct.

The mediation analysis shows that the indirect effect via Ethical Alignment was not statistically significant (Estimate=0.179, SE=0.159, $z=1.124$, $p=.261$, 95% CI [-0.118, 0.517]). The path from group to Ethical Alignment was also not significant (Estimate=0.266, SE=0.240, $z=1.110$, $p=.267$, 95% CI [-0.181, 0.749]), indicating that the experimental manipulation did not significantly elevate participants' perceived sense of moral alignment with the platform's recommendations. As with platform trust, Ethical Alignment was a strongly significant predictor of Satisfaction (Estimate=0.673, SE=0.066, $z=10.23$, $p<.001$, 95% CI [0.544, 0.800]), demonstrating its general psychological importance. However, it did not function as a mediating mechanism in this experimental context (see full mediation output in Figure C8.9). **H4 is not supported.** As with Platform Trust, the convergence of Ethical Alignment, Platform Trust, and Satisfaction onto a single EFA factor (Section 4.1.2) indicated that these constructs may not be empirically separable in a brief experimental context, which limits the interpretive confidence of this mediation test. When Ethical Alignment entered the model as a mediator, a marginally significant negative effect of group condition on satisfaction emerged (Estimate= -0.320, SE= 0.168, $z=-1.901$, $p=.057$, 95% CI [-0.660, 0.000]). Although the direction of this effect resembles the pattern observed in the H2 model, the indirect effect through Ethical Alignment was non-significant ($p=.261$). Accordingly, the results do not provide evidence of a suppressor effect comparable to that observed in H2.

4.4.5 H5: Moderation by Algorithmic Literacy

Because the sample size constrained the moderation test, H5 was evaluated in two stages. First, a confirmatory test was conducted using PROCESS Model 1 (Hayes, 2022) in JASP, examining a single interaction term (Group x AL) predicting satisfaction directly. Second, three exploratory moderated mediation models (PROCESS Model 7, R version 4.4.1) were estimated to test whether algorithmic literacy moderated the path from ethical customisation to one of the three organism-level variables (perceived agency, platform trust, and perceived ethical alignment). The non-significance reported below should therefore be interpreted cautiously, as neither approach had sufficient power to rule out conditional effects at the individual mediated pathway level.

A significant baseline imbalance in Algorithmic Literacy was observed between groups despite random assignment. Treatment-group participants reported lower Algorithmic Literacy ($M=3.745$, $SD= 1.431$) than control-group participants ($M=4.293$, $SD=1.439$): $t(120) = 2.108$, $p=.037$, $d=0.382$. This moderate imbalance reduces interpretive confidence for the moderation test and underscores the importance of treating Algorithmic Literacy as a covariate throughout all models.

The confirmatory moderated regression (PROCESS Model 1) tested the interaction term Group_d x AL_comp as a predictor of satisfaction with all covariates included. The interaction term was non-significant (Estimate = 0.073, SE = 0.157, $z=-0.462$, $p=.644$, 95% CI [-0.235, 0.381]), providing no

support for moderation at any percentile of Algorithmic Literacy. The main effect of AL_comp on Satisfaction remained significant and positive (Estimate = 0.296, $p < .001$; Table 8). **H5 is therefore not supported.**

Table 8

Moderation Analysis: Group_d * Algorithmic Literacy Interaction Predicting Satisfaction

Path	Estimate	SE	z	p	95% CI Lower	95% CI Upper
Group_d → S_comp	0.007	0.228	0.029	.977	-0.440	0.453
AL_comp → S_comp	0.296	0.078	3.773	< .001	0.142	0.450
Group_d × AL_comp → S_comp	0.073	0.157	0.462	.644	-0.235	0.381

Note. N = 122. Moderation effect estimates based on mean-centred variables. $R^2 = 0.109$. The interaction term Group_d × AL_comp did not reach significance, indicating no moderation of the group–satisfaction path by Algorithmic Literacy.

Source: Author's own elaboration.

To more faithfully operationalise H5 as originally stated, three exploratory moderated mediation models (PROCESS Model 7, Hayes, 2022) were also estimated in R (version 4.4.1), each examining whether Algorithmic Literacy moderated the path from experimental group to one organism-level mediator. All three indices of moderated mediation had confidence intervals that included zero, thereby confirming the absence of a significant moderated mediation across all three pathways. The most theoretically coherent pattern emerged in the perceived agency model. The indirect effect of ethical customisation on satisfaction through perceived agency was significant at the 16th percentile of algorithmic literacy (Effect=0.561, 95% CI [0.175, 1.007]) and the 50th percentile (Effect=0.389, 95% CI [0.121, 0.713]), but non-significant at the 84th percentile (Effect= 0.289, 95% CI [-0.061, 0.701]), suggesting that visible ethical controls may be more agency-activating for users with less algorithmic knowledge. For platform trust, conditional indirect effects shifted from a near-zero positive value at low literacy (Effect= 0.071, 95% CI [-0.494, 0.569]) to increasingly negative values at higher literacy levels (Effect= -0.090 at the median; Effect= -0.183 at the 84th percentile), however none reached significance. This directional pattern is consistent with a scepticism account in which more algorithmic literate users are less reassured by surface-level ethical interface signals.

Table 9

Exploratory Moderated Mediation: Path Interactions and Conditional Indirect Effects (PROCESS Model 7)

	Model A Perceived Agency	Model B Platform Trust	Model C Ethical Alignment
Path interaction: Group × AL → Mediator			
Estimate	-0.188	-0.112	-0.062
SE	0.168	0.159	0.167

p	.267	.481	.711
ΔR^2 (X×W)	.009	.004	.001
Conditional indirect effects: Group → Mediator → Satisfaction			
AL at 16 th percentile	0.561 [0.175, 1.007] *	0.071 [-0.494, 0.569]	0.394 [-0.175, 0.938]
AL at 50 th percentile	0.389 [0.121, 0.713] *	-0.090 [-0.456, 0.296]	0.310 [-0.056, 0.715]
AL at 84 th percentile	0.289 [-0.061, 0.701]	-0.183 [-0.701, 0.353]	0.261 [-0.211, 0.763]
Index of moderated mediation	-0.086 [-0.255, 0.089]	-0.080 [-0.308, 0.175]	-0.042 [-0.264, 0.186]

*Note. N = 121 (one case excluded due to missing data). PROCESS Model 7 (Hayes, 2022), estimated in R (version 4.4.1) using the official PROCESS macro for R. Bootstrap resamples = 5,000, seed = 1234. Algorithmic Literacy (AL) is mean-centred. Values in brackets are bootstrapped 95% confidence intervals. * Indicates CI excludes zero (significant conditional indirect effect). ΔR^2 reflects the unique variance explained by the X×W interaction in the mediator model. All models include eco-ethical values, manipulation check awareness, and price items as covariates. Full moderated mediation output for Models A, B, and C is presented in Figures C8.11–C8.13.*

Source: Author's own elaboration.

4.4.6 Summary of Hypothesis Testing

Of the five proposed hypotheses, one received support (H2) and four were not supported (H1, H3, H4, H5). These findings introduce important theoretical and practical implications, which are discussed in Chapter 5.

Table 10
Summary of Hypothesis Testing Outcomes

H	Hypothesis	Test Used	Outcome
H1	Perceived availability of ethical customisation → Satisfaction (direct)	Independent samples t-test; Hierarchical regression M ₁ (JASP)	Not supported
H2	Ethical customisation → Perceived Agency → Satisfaction (mediation)	Maximum likelihood mediation (JASP)	Supported ✓
H3	Ethical customisation → Platform Trust → Satisfaction (mediation)	Maximum likelihood mediation (JASP)	Not supported
H4	Ethical customisation → Ethical Alignment → Satisfaction (mediation)	Maximum likelihood mediation (JASP)	Not supported
H5	Algorithmic Literacy moderates customisation → organism variables	PROCESS Model 1 (confirmatory, JASP);	Not supported (confirmatory)

		PROCESS Model 7 (exploratory, R)	Directionally meaningful (exploratory)
--	--	-------------------------------------	--

Source: Author's own elaboration.

4.5 Exploratory Analysis: Behavioural Intention

As an exploratory supplement to the primary hypotheses, the item “I could imagine myself using a booking platform that includes an AI-powered recommendation system like this in a real booking situation” (BI1_r) was examined separately as a single-item indicator of behavioural intention.

The item correlated strongly with the three-item Satisfaction composite ($r=0.697$, $p<.001$), as well as with Ethical Alignment ($r=0.553$, $p<.001$) and Platform Trust ($r=0.537$, $p<.001$), suggesting that the psychological antecedents of behavioural intention in this context are individual-level evaluative constructs rather than experimental conditions. While satisfaction was the strongest predictor of whether users would consider using the platform in a real booking situation, the positive relationship with perceived ethical alignment and platform trust also indicates that users are more likely to engage with the platform when they see it as trustworthy and consistent with their own values. These findings are consistent with the broader patterns of results and are discussed further in Section 5.6.

4.6 Qualitative Comments: Supplementary Insights

Although the study employed a quantitative methodology, 15 participants (6.3% of the original sample) provided optional free-text responses at the end of the survey. While the limited number of responses prevents a systematic qualitative analysis, several comments offer theoretically meaningful supplementary insights that contextualise the quantitative findings.

Most notably, a retained control-group participant explicitly references their extensive booking history as a mechanism through which they expected platform recommendations to already be personalised, stating that “a platform on which I have already made 100+ bookings knows much more about me and makes better suggestions.” This comment provides direct qualitative evidence for the manipulation check ambiguity discussed in the limitations section: the respondent’s understanding of implicit algorithmic influence appears to have coexisted with a correct response to the manipulation check item for the non-manipulation within the control group, suggesting that the conceptual boundary between passive personalisation and active ethical customisation may remain ambiguous even for engaged and experienced users.

A second substantively rich comment came from a retained treatment-group participant who, despite having perceived the ethical customisation features, expressed deep scepticism about whether platform-level ethical controls translate into genuine recommendation changes: “On all booking platforms, it is about business, meaning making money. Therefore, suggestions are primarily generated according to this criterion.” This comment articulates a form of ethics-washing scepticism that is consistent with the negative correlation observed between ethical consumption orientation and satisfaction in the quantitative data. It suggests that for some users, the visibility of ethical features may be perceived as performative rather than substantive, which diminishes rather than amplifies the positive psychological effects the study sought to measure.

Two excluded treatment-group participants expressed strong AI anxiety, suggesting that emotional responses to AI may hinder attentive processing of interface features and contribute to manipulation check failure independently of algorithmic literacy.

Taken together, these comments do not alter the quantitative conclusions of the study but provide a layer of validity to several of its theoretical claims, particularly regarding algorithmic scepticism, the

blurred boundary between personalisation and customisation, and the emotional dimensions of AI trust that structured Likert scales may insufficiently capture.

Chapter 5: Discussion

This study aimed to investigate whether granting travellers the perceived ability to customise the ethical parameters of an AI-powered booking platform would enhance their satisfaction (specifically through the visible presence of ethical preference sliders on a Booking.com interface), and through which psychological pathways this effect would operate. Drawing on the S-O-R framework, this study tested five hypotheses relating to direct effects, mediated pathways via perceived agency, platform trust, and ethical alignment, and the moderating role of algorithmic literacy. This chapter discusses the theoretical and practical significance of the findings in light of the existing literature.

5.1 The Absence of a Direct Effect: H1

Ethical customisation did not directly increase traveller satisfaction, suggesting that its influence operates indirectly through psychological mechanisms rather than through immediate evaluative responses. This finding indicates that the visibility of ethical controls may be insufficient unless users receive those controls as meaningful and consequential.

At the same time, the regression analysis points to a broader and more complex explanation of traveller satisfaction. Although the experimental condition itself did not significantly predict satisfaction, the full regression model (including algorithmic literacy, eco-ethical values, and price perceptions) was statistically significant, explaining 15.2% of the variance in traveller satisfaction.

5.2 Perceived Agency as a Partial Mediator: H2

The confirmation of perceived agency as a significant mediator (H2) represents the central theoretical finding of this thesis. The mediation result points to something more specific, where users did not respond to the sliders as a visual feature. Instead, they responded to the sense of personal control and moral self-determination it signalled.

When perceived agency is not in the model, the treatment group shows no significant satisfaction advantage over the control group. Once Perceived Agency is included, a positive indirect path through Perceived Agency and a counteracting negative direct path both emerge. This means that ethical customisation simultaneously boosts satisfaction through felt control and slightly suppresses it through raised awareness of data practices. The net result is a null total effect that conceals two competing psychological processes.

This outcome is in line with the Self-Determination Theory, which suggests that autonomy is a basic psychological need and that circumstances that let people make their own decisions typically result in more positive assessments and higher levels of intrinsic motivation (Ryan & Deci, 2020). In the context of AI-powered travel recommendations, where algorithmic systems are often experienced as opaque or difficult to understand, the availability of visible ethical controls seems to shift the perceived locus of decision-making back toward the user. Even participants who only viewed a static image (without actively using the sliders) perceived greater control, which was sufficient to elevate Perceived Agency scores.

The findings also resonate with Signalling Theory as introduced by Michael Spence (1973). The ethical customisation sliders appear to function as a signal that the platform values user governance and transparency. Importantly, this signal operates before any actual interaction occurs, suggesting that interface design itself can shape psychological expectations about autonomy and control.

The negative and statistically significant remaining direct effect once perceived agency is controlled is a theoretically important finding that warrants explicit discussion. Two mechanisms may jointly explain this. First, visible ethical sliders introduce decision-making responsibility absent in the control condition, consistent with choice overload theory (Long et al., 2025). Second, this pattern aligns with the privacy paradox (Zhang & Sundar, 2019). Making privacy controls visible can simultaneously

reassure users and heighten awareness of the very data practices those controls address. Together, these two mechanisms suggest that ethical customisation is a double-edged interface feature. It simultaneously empowers and unsettles. Its net effect on satisfaction is positive only because the agency-activation pathway outweighs the friction, cognitive burden, and raised data-awareness it also introduces.

This finding has direct theoretical implications for the S-O-R framework. The organism-level response to the stimulus is not uniformly positive but involves competing psychological processes, of which perceived agency is the dominant one in single-exposure contexts. The present results extend Yang et al.'s (2024) findings on algorithmic fatigue by demonstrating that even symbolic forms of control (visible sliders that were never actively used) are sufficient to partially restore felt autonomy, while simultaneously making the very data practices they are intended to address more prominent.

5.3 Platform Trust as a Non-Mediator: H3

Hypothesis 3, which proposed that platform trust would mediate the relationship between ethical customisation and traveller satisfaction, was not supported. Participants in the treatment condition did not report significantly higher levels of platform trust compared to the control group, and the indirect effect through trust was also non-significant. This result is somewhat unexpected given the substantial body of literature arguing that transparency mechanisms and XAI features are key drivers of user trust in digital platforms. Therefore, trust is a construct that is built through repeated use, not a single interface image (Afroogh et al., 2024).

A second explanation concerns the direction of the effect itself. The relationship between the treatment condition and platform trust was slightly negative, which (while non-significant) runs contrary to the hypothesised direction. This pattern is theoretically coherent under two complementary accounts.

Under a transparency backfire account (consistent with Zhang & Sundar, 2019), making algorithmic configurability visible simultaneously communicates that the platform acknowledges its own opacity. Although trust did not mediate the experimental effect, the directionally lower Platform Trust scores in the treatment group suggest that the visibility of ethical controls may have heightened participants' awareness of the underlying algorithmic process rather than reassuring them of its integrity. This is consistent with the privacy paradox identified by Zhang & Sundar's (2019), whereby increased transparency about a system's configurability can simultaneously raise users' consciousness of how much control they ordinarily lack, amplifying scepticism rather than confidence.

Under a credibility gap account, users may perceive the ethical sliders as interface signals without being able to verify that they connect to actual recommendation logic, which produces scepticism rather than reassurance. Both accounts predict that trust requires not only the signalling of ethical intent but evidence of its implementation. This interpretation is supported by qualitative data, in which a treatment-group participant explicitly articulated doubt that platform-level ethical controls translate into genuine recommendation changes.

5.4 Perceived Ethical Alignment as a Non-Mediator: H4

Hypothesis 4, which proposed that perceived ethical alignment would mediate the relationship between ethical customisation and traveller satisfaction, was not supported. The experimental manipulation did not significantly increase participants' perceptions that the AI-generated recommendations reflected their personal values, and the indirect effect through ethical alignment was non-significant as well. Perceived Ethical Alignment is a construct that this thesis introduces. The null mediation result, therefore, does not invalidate the construct, but it specifies the conditions under which it activates.

Unlike perceived agency, which responds to the visible presence of user controls, ethical alignment requires evidence that the system's outputs genuinely reflect the user's moral priorities, which is a condition that a static, pre-configured interface cannot fulfil without active interaction and observable recommendation change. Instead, users may need evidence that the platforms meaningfully incorporate their ethical priorities into the recommendation process itself. Future experimental designs that allow participants to actively manipulate ethical settings and observe personalised recommendation changes may be better suited to capturing this process.

The pattern of a strong within-sample correlation alongside the negligible between-group difference constitutes indirect but meaningful evidence for a distinction between two modes of ethical alignment's operation. As a within-person evaluative construct, it is a powerful predictor of satisfaction. As a between-person response to a single interface exposure, it is resistant to experimental manipulation. This distinction aligns with a broader theoretical point, where perceived ethical alignment may function less like a momentary impression and more like a considered moral judgement, which requires interactive evidence. This positions it theoretically closer to value congruence (a judgement that accumulates through experience) than to a first-impression perceptual response. Future research could test this by positioning Ethical Alignment as a moderator of the customisation-satisfaction relationship rather than a mediator, or by using longitudinal designs in which value-alignment perceptions are measured across multiple platform interactions. The present study's contribution is to establish the boundary conditions under which perceived ethical alignment does and does not function as an experimentally activatable state, which is itself a theoretical step in the construct's development.

5.5 The Role of Algorithmic Literacy: H5

The hypothesised moderating role of algorithmic literacy on the relationship between ethical customisation and the organism-level responses (H5) was not supported. The confirmatory interaction term (Group x AL, PROCESS Model 1) was non-significant, and conditional effects at all literacy levels were likewise non-significant. At the same time, algorithmic literacy emerged as a significant positive predictor of traveller satisfaction in the regression analysis, independently explaining a notable portion of the variance. This suggests that while literacy did not alter how participants responded to the experimental manipulation itself, it still played an important role in shaping overall attitudes toward AI-powered recommendations regardless of condition.

The absence of a moderation effect aligns with recent critiques in the algorithmic literacy literature, which argue that the conditions under which literacy promotes or inhibits critical engagement with algorithmic systems remain highly context-dependent and underexplored (Gagrčin et al., 2024). The exploratory moderated mediation analyses (PROCESS Model 7) provided a more granular but similar non-significant picture that was consistent with the confirmatory result. All three indices of moderated mediation had confidence intervals that included zero, confirming the absence of significant moderated mediation. However, conditional indirect effects within the perceived agency model revealed a theoretically meaningful directional pattern. The agency-mediated effect was significant at both the 16th percentile and the 50th percentile of algorithmic literacy but diminished to non-significance at the 84th percentile. This pattern suggests that the agency-activating effect of ethical customisation may be stronger for users with lower algorithmic literacy, while more literate users apply stricter evaluative standards that weaken this effect. For platform trust, conditional indirect effects were negative and grew more negative at higher literacy levels, consistent with a scepticism account in which knowledgeable users are less persuaded by surface-level ethical interface signals (Bietti, 2021). These directional patterns are consistent with the broader argument that algorithmic literacy operates differently across distinct psychological pathways, a nuance that aggregate moderation tests on the outcome variable alone cannot detect.

The direct positive relationship between Algorithmic Literacy and Satisfaction also highlights an important implication. Users with a stronger understanding of algorithms may feel more confident

and comfortable when interacting with AI-powered recommendation systems. Because they have a better understanding of how these systems work, including their strengths and limitations, they may be more likely to appreciate and benefit from AI-powered recommendations. Users who understand how recommendations work seem to get more out of them, regardless of the interface. Investment in user education may be as commercially valuable as investment in better algorithms.

5.6 Synthesis and Theoretical Contributions

Taken together, the findings of this study make a clear and theoretically coherent contribution to the literature. Within the S-O-R framework, the perceived availability of ethical customisation features operates as an effective stimulus, but only in a constrained sense. Its primary psychological effect is the elevation of Perceived Agency, which in turn positively influences satisfaction. The results suggest that ethical customisation operates primarily as an agentic experience in single-exposure contexts, and that its value to the traveller lies not only in the tangible output of ethically filtered recommendations, but in the psychological reassurance that they retain authorship over the moral logic of the AI.

A further finding calling for theoretical discussion is the negative correlation between eco-ethical values orientation and traveller satisfaction, observed consistently across all model specifications as a significant negative predictor. This counterintuitive pattern is interpretable through the EDP (Oliver, 1980; Oliveri et al., 2019). Users who hold stronger pre-existing ethical values enter the interaction with higher moral expectations of the platform. A static booking interface is unlikely to fully satisfy these expectations, producing negative disconfirmation precisely among users most likely to value ethical customisation. This suggests that the current experimental stimulus, while sufficient to activate perceived agency in the treatment group, may not have been strong enough to meet the elevated evaluative standards of ethically oriented users (the practical implications of this pattern are further discussed in Section 6.2).

The non-confirmation of trust and ethical alignment as mediators contributes equally important theoretical implications. It suggests that the psychological impact of ethical interface features is more immediate and agency-centric than value-resonant in single-exposure models. Long-term, repeated interactions may be necessary to cultivate trust and value coherence, while a sense of agency appears to be more rapidly activated by the visual availability of control-conferring features.

A noteworthy pattern reinforces this distinction. Ethical Alignment's negligible between-group difference stands in strong contrast to its exceptionally strong within-sample correlation with satisfaction. This reveals a fundamental asymmetry between the two primary organism-level constructs. Perceived agency can be activated rapidly through the visible presence of user controls, demonstrating that the signal of configurability alone is sufficient. On the contrary, ethical alignment operates as a necessary but not sufficient condition. When users genuinely experience it, it powerfully shapes satisfaction, but a single static exposure to preconfigured sliders cannot produce that experience. It requires sustained, interactive engagement in which users can directly observe that the system's outputs shift meaningfully in response to their moral preferences. Future research should therefore conceptualise ethical alignment not as a short-term manipulable state, but as a stable evaluative judgement that emerges only through demonstrably value-responsive system behaviour.

It is important to acknowledge that the convergence of Ethical Alignment, Platform Trust, and Satisfaction onto a single EFA factor points to a limitation in discriminant validity. This does not invalidate the mediation analyses for H3 and H4, since these were estimated correctly as separate structural models. However, it does suggest that the non-significant mediation effects should be interpreted with caution. Rather than concluding that trust and ethical alignment are unimportant for satisfaction, the findings more likely indicate that these constructs are difficult to distinguish empirically within a single-exposure experimental setting. This supports the broader theoretical

argument of the thesis that trust and ethical alignment may require repeated interactions and longer-term engagement with the platform before they emerge as clearly independent psychological constructs.

More broadly, the findings reframe the value of ethical AI interface features in tourism. Their significance lies not in their signalling function alone (communicating that a platform has ethical commitments), but in their capacity to reposition the traveller from a passive recipient of algorithmic logic to a perceived co-author of their own booking experience. This repositioning is psychological before it is behavioural. The present study demonstrates that it occurs even when users never interact with the controls at all.

Chapter 6: Conclusions

6.1 Summary

This thesis investigated whether the perceived availability of ethical customisation features in an AI-powered travel booking platform affects traveller satisfaction, and through which psychological mechanisms. Using a between-subjects experimental design (N=122) grounded in the S-O-R framework, participants were randomly assigned to view either a standard Booking.com interface or an enhanced version featuring visible ethical preference sliders related to eco-friendliness, data privacy, and social impact. The study further examined perceived agency, platform trust, and perceived ethical alignment as mediating mechanisms, while algorithmic literacy was tested as a potential moderator.

Of the five hypotheses tested, only H2 received empirical support. Ethical customisation features improved satisfaction indirectly by strengthening users' perceived agency, rather than through a direct evaluative response.

Overall, the findings suggest that ethical customisation features do not automatically enhance satisfaction at first exposure, but they do meaningfully influence how users psychologically experience AI-powered recommendation systems. The results highlight Perceived Agency as the key mechanism through which ethical interface design shapes user responses. Visible opportunities for user input and control appear to matter more than ethical signalling alone, particularly in AI-mediated decision environments where users may otherwise feel detached from or uncertain about algorithmic processes. Qualitative comments from participants reinforce these patterns, with control-group participants expressing scepticism toward opaque recommendation processes and treatment-group participants acknowledging both the reassurance and the data-awareness that ethical controls introduced.

The central empirical contribution is that of the three proposed psychological pathways, only perceived agency transmitted the effect of ethical customisation on satisfaction. Rather than enhancing satisfaction through ethical branding or signalling alone, customisation features appear most effective when they restore users' felt authorship over AI-generated outcomes, with direct implications for how responsible AI should be designed, communicated, and evaluated in digital travel contexts.

6.2 Managerial implications

The findings of this thesis carry several concrete implications for product managers, UX designers, and platform strategists working in AI-powered travel and hospitality contexts. The most direct managerial insight is that Perceived Agency is a meaningful pathway to Traveller Satisfaction, and that it can be reliably activated through the visible availability of ethical customisation features. This does not require users to actively engage with sliders or settings. The perception that such options exist is itself sufficient to shift psychological experience. For practitioners, this implies that investing in the communicative design of ethical controls is as important as their technical implementation. Platforms should ensure that customisation features are predominantly visible, clearly labelled, and accompanied by plain-language explanations, as demonstrated by the tooltip descriptions in this study's treatment interface.

Second, the unexpected finding that platform trust was directionally lower in the treatment group deserves managerial attention. While not statistically significant, this pattern suggests that introducing visible algorithmic controls may raise the salience of the AI's underlying decision-making in ways that activate scepticism rather than reassurance, at least in a single-exposure context. This underscores that ethical customisation features must be accompanied by genuine transparency about how preferences translate into recommendation outcomes. If users cannot observe a clear

and credible link between their ethical choices and the results they receive, the promise of control risks being perceived as performative. Platforms should therefore invest not only in visible controls but in feedback mechanisms that make the downstream effect of ethical preferences visible to users in real time.

Third, the strong and consistent positive effect of Algorithmic Literacy on Traveller Satisfaction across all models suggests that platforms would benefit from investing in user education. Onboarding initiatives, contextual tooltips, and brief explainers that help users understand how AI recommendations are generated could raise effective algorithmic literacy and thereby amplify the positive effects of ethical customisation. This is especially relevant given the relatively low mean algorithmic literacy score observed in this sample. Bridging this knowledge gap represents a practical opportunity for platforms seeking to enhance user experience and long-term trust.

A fourth implication concerns the counterintuitive negative relationship between Eco-ethical Value orientation and Satisfaction, which requires strategic consideration. As discussed in Section 5.6, this pattern is interpretable through the EDP. Users with stronger ethical values likely entered the interaction with higher moral expectations, which a static interface may not have fully satisfied. As a result, the users who cared most about ethical considerations were also the most likely to experience disappointment. For platform designers, this has important practical implications. Ethically conscious users may represent both the greatest reputational challenge and the greatest opportunity. If platforms fail to meet their expectations, dissatisfaction may increase. However, genuinely effective ethical customisation features, where user preferences clearly and transparently influence recommendation outcomes, could strengthen trust and improve satisfaction among these highly value-driven users.

6.3 Theoretical implications

This study makes theoretical advances along three dimensions. First, it extends the S-O-R framework by introducing what can be described as a perceptual ethical stimulus, namely, the visible availability of customisation features. While previous applications of the S-O-R model in e-commerce and tourism have focused on stimuli such as website quality, pricing, or review valence, this study shows that ethical interface design can function as an equally meaningful stimulus category.

Second, it builds on Zhang & Sundar's (2019) distinction between personalisation and customisation by extending it into the ethical domain. In doing so, it provides experimental evidence that giving users control over moral parameters produces psychological effects that are not simply extensions of functional control but are conceptually distinct.

Third, it contributes to the emerging literature on AI ethics in tourism, which has so far been largely theoretical or based on survey research, by offering causal experimental evidence. This addresses a recognised gap in the field, where stronger experimental designs have been called for but remain relatively rare.

6.4 Limitations and implications for future research

The present study has several limitations that should be acknowledged. First, the static interface stimulus limits findings to perceptual availability effects rather than behavioural engagement, making replication with interactive prototypes necessary (Aguinis & Bradley, 2014).

Second, although the sample size was sufficient to detect medium effects in a two-group between-subjects design, it is less suitable for more complex analyses such as moderated mediation. In particular, the test of H5 was constrained by limited variance in algorithmic literacy and small subgroup sizes, meaning that the moderation results should be interpreted as exploratory.

The instructed-response attention item was accidentally dropped during the final revision. This was addressed by applying time-based exclusion instead, though it is recognised as an imperfect substitute and therefore a limitation of this thesis.

Third, recruitment through social media and academic networks produced a convenience sample that likely overrepresents digitally literate and ethically conscious individuals. This may bias responses toward more favourable evaluations of ethical customisation features. Social desirability bias may also have influenced responses, as visibly ethical interface elements could encourage participants to provide socially acceptable responses independently of genuine psychological effects. Additionally, while the survey was administered in four languages following review for semantic equivalence, the absence of formal back-translation procedures represents a validity threat for constructs with abstract ethical content, particularly perceived ethical alignment, where translation may differ meaningfully across cultures. Furthermore, the significant between-group imbalance in algorithmic literacy, despite random assignment, represents a residual threat to internal validity for H5 that statistically covariate control can help to reduce but not fully eliminate.

Another limitation lies in the potential for common method bias, since all psychological variables were collected in the same survey session from the same participants using self-reporting items (Podsakoff et al., 2003). Although the experimental group assignment was objectively manipulated and therefore less affected by bias, relationships between mediators and outcomes may still have been inflated by shared measurement methods.

Three directions for future research emerge directly from these limitations. First, studies using interactive interfaces could examine whether active engagement with ethical controls strengthens the agency-activation effect. Second, larger multi-country samples combined with formal back-translation and measurement invariance testing could explore cross-cultural differences in ethical preference formation. Finally, longitudinal designs could determine whether the effects of ethical customisation persist over time or diminish as novelty effects fade.

Appendices

Appendix A: Survey Instrument (English Version)

Introduction and consent

Welcome to this survey!

This study examines how users perceive AI-powered recommendation systems on booking platforms.

You will be shown a sample booking interface with hotel suggestions. Please focus your assessment on the platform and the recommendation system – not on the individual hotels.

In this survey, you will be shown one of several possible user interfaces at random; these may differ slightly in terms of individual functions or visual elements.

The survey will take approximately 15 minutes. Your answers are anonymous and will be used exclusively for academic research purposes.

By continuing, you confirm that your participation is voluntary.

Thank you very much for taking part!

This survey is anonymous.

The record of your survey responses does not contain any identifying information about you, unless a specific survey question explicitly asked for it.

If you used an identifying access code to access this survey, please rest assured that this code will not be stored together with your responses. It is managed in a separate database and will only be updated to indicate whether you did (or did not) complete this survey. There is no way of matching identification access codes with survey responses.

[Standard LimeSurvey anonymity notice, generated automatically by the survey platform]

Section 1- Experience

Q1. How often have you used online booking platforms (e.g., Booking.com, Airbnb, Expedia) in the past 24 months?

Never

One to two times

3 to 5 times

More than 5 times

Section 2- Ethical orientation (Pre-exposure)

Q2. Please indicate to what extent you agree with the following statements:

1 = Strongly Disagree

2

3

4

5

6

7 = Strongly Agree

- It is important to me that travel options are environmentally friendly.
- It is important to me that travel options support local communities.
- It is important to me that my personal data is handled responsibly.

Section 3- Scenario and Stimulus [see Appendix B]

Section 4- Manipulation check (both groups)

Q3. In the version I saw, it appeared that users could adjust how the recommendations were generated.

Yes

No

Section 5- Perceived Agency

This section focuses on the sense of control you had regarding the shown picture. Please indicate your level of agreement with the following statements.

Q4. Considering the AI-powered recommendation system on this booking platform:

[7-point Likert: 1 = Strongly Disagree to 7 = Strongly Agree]

- It appears to give users the ability to influence how the AI-powered recommendations are generated.
- It appears to allow users to be in control of the recommendations.
- It would allow me to effectively express my preferences.

Section 6- Perceived Ethical Alignment

This section focuses on how well the AI-powered recommendations align with your personal values. Please indicate your level of agreement with the following statements.

Q5. Considering the recommendations generated by the AI-powered recommendation system on this booking platform:

[7-point Likert: 1 = Strongly Disagree to 7 = Strongly Agree]

- They seem to match my personal values.
- They seem to reflect what matters to me when choosing travel options.
- They seem to take into account aspects that are important to me personally.

Section 7- Platform Trust

This section is about how much you trust the recommendation system. Please indicate your level of agreement with the following statements.

Q6. Considering the AI-powered recommendation system on this platform:

[7-point Likert: 1 = Strongly Disagree to 7 = Strongly Agree]

- I trust it to provide reliable recommendations.
- I believe it acts in my best interest.
- It operates in a transparent way.

Section 8- Satisfaction

This section focuses on your overall evaluation of the booking platform. Please indicate your level of agreement with the following statements.

[7-point Likert: 1 = Strongly Disagree to 7 = Strongly Agree]

- Q7. I would be satisfied with the AI-powered recommendations on this booking platform.
- Q8. The AI-powered recommendations on this booking platform seem to meet my expectations.
- Q9. Using this booking platform, supported by its AI-powered recommendation system, provides a pleasant user experience.
- Q10. I could imagine myself using a booking platform that includes an AI-powered recommendation system like this in a real booking situation.

Section 9- Algorithmic Literacy

This section focuses on your familiarity with and understanding of AI-powered recommendation systems. Please indicate your level of agreement with the following statements.

[7-point Likert: 1 = Strongly Disagree to 7 = Strongly Agree]

- Q11. I am familiar with how AI-powered recommendation systems personalise suggestions on booking platforms.
- Q12. I feel confident in my understanding of how AI-powered recommendation systems work on booking platforms.
- Q13. I have a general understanding of how my data is used by AI-powered recommendation systems on booking platforms.

Section 10- Price Sensitivity

This section focuses on your preferences when making booking decisions. Please indicate your level of agreement with the following statements.

[7-point Likert: 1 = Strongly Disagree to 7 = Strongly Agree]

- Q14. I prioritise ethical aspects over price when booking.
- Q15. I am willing to pay a higher price for a hotel that meets my personal expectations.

Section 11- Demographics

Finally, this last section allows us to categorise the results statistically and anonymously. These details are essential to ensure the scientific validity of this study and will be used strictly for academic research purposes.

Q16. How old are you?

- Under 18
- 18-24
- 25-34

35-44
45-54
55-64
65 and older

Q17. What is your gender?

Male
Female
Prefer not to say
Other

Q18. Where are you from?

Belgium
Netherlands
France
Luxembourg
Germany
Other

Q19. What is your current employment status?

Student
Employed full-time
Employed part-time
Self-employed / Freelancer
Unemployed / Looking for work
Retired
Unable to work
Prefer not to say

Section 12- Closing thoughts (optional free text)

You can use this section to share any additional thoughts or opinions.
This question is optional, but your feedback is greatly appreciated.

Section 13- Conclusion

Thank you for participating!
Your responses will be a tremendous help as I work toward earning my degree at HEC Liège!
If you have any questions about this study, please feel free to contact me at the following address:
Jana.Gennen@student.uliege.be.

Appendix B: Experimental Interface Stimuli

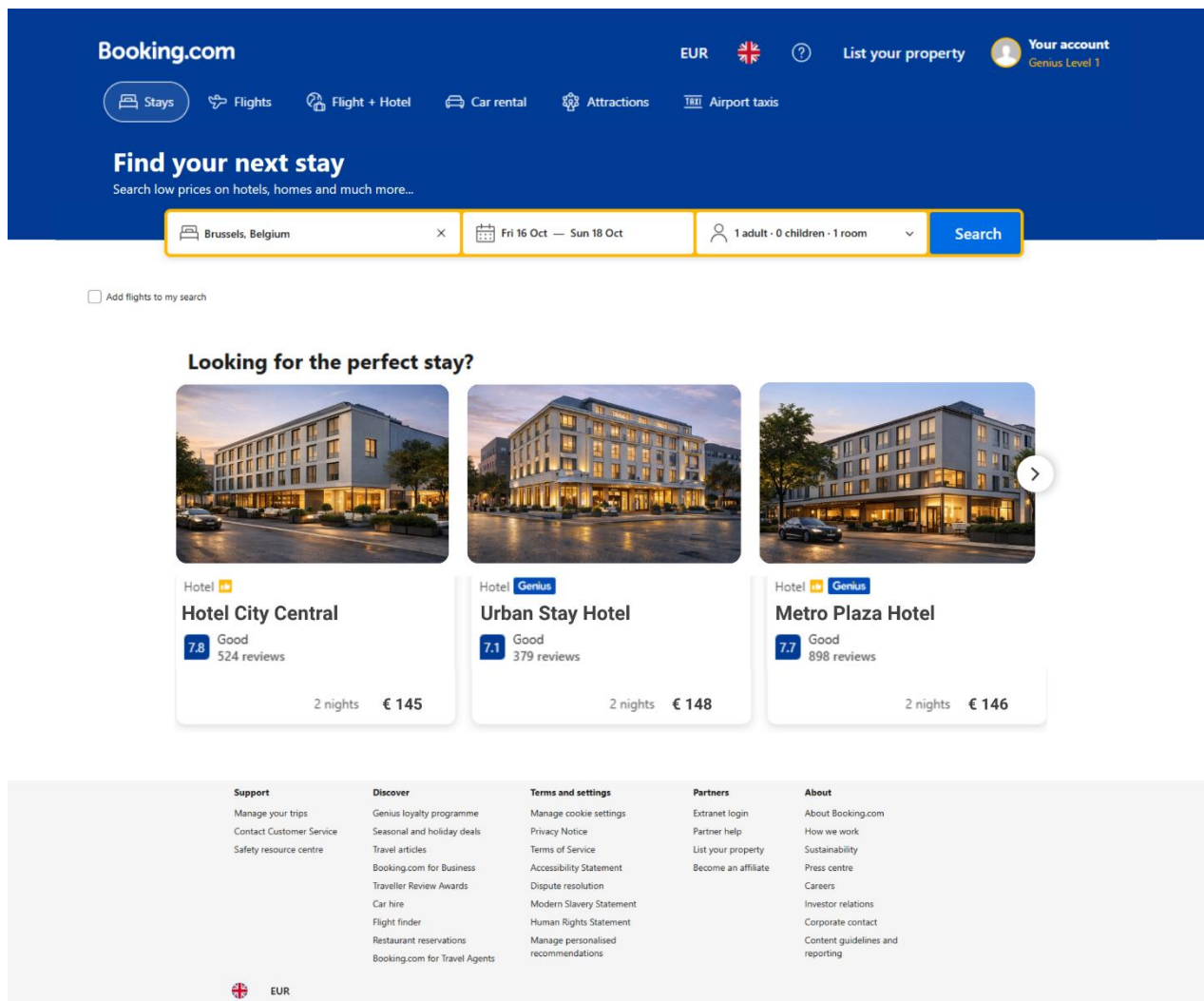


Figure B1. Control condition interface (Group A): Standard Booking.com hotel search results displaying three hotels (Hotel City Central, €145, Urban Stay Hotel, €148, Metro Plaza Hotel, €146) without ethical customisation features. Source: Author's own design, adapted from Booking.com interface.

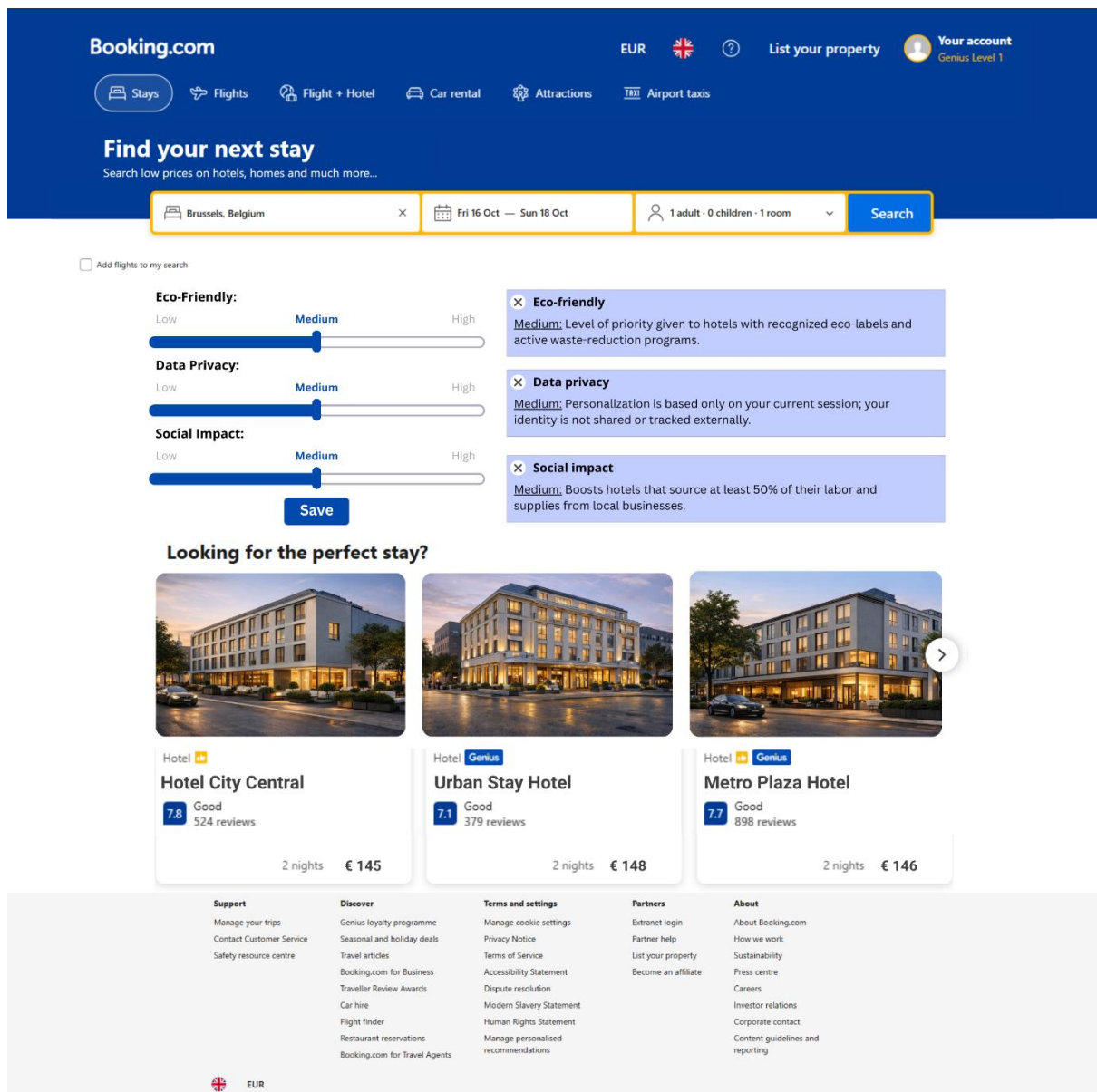


Figure B2. Treatment condition interface (Group B): Booking.com interface with three ethical preference sliders (Eco-Friendly, Data Privacy, Social Impact), each pre-set to Medium and accompanied by explanatory tooltips. Source: Author's own design, adapted from Booking.com interface.

Appendix C: Supplementary Statistical Output

Appendix C presents the complete statistical output for all analyses reported in Chapter 4. All analyses were conducted in JASP (version 0.97). Output is presented as screenshots taken directly from the software and is organised in the order in which results are reported in the main text.

C1 Reliability statistics for all scales (Cronbach's α , item-dropped statistics)

1. Perceived Price (PS)

Frequentist Scale Reliability Statistics

Coefficient	Estimate	Std. Error	95% CI	
			Lower	Upper
McDonald's ω				
Cronbach's α	0.336			

Note. Omega calculation with CFA failed. Try changing to PFA in 'Advanced Options'. The analytic confidence interval is not available for coefficient alpha/lambda2 when data contain missings and pairwise complete observations are used. Try changing to 'Delete listwise' within 'Advanced Options'.

Frequentist Individual Item Reliability Statistics

Item	McDonald's ω (if item dropped)			Cronbach's α (if item dropped)		
	Estimate	Lower 95% CI	Upper 95% CI	Estimate	Lower 95% CI	Upper 95% CI
PS1_r						
PS2_r						

Note. Please enter at least 3 variables for the if item dropped statistics.

Figure C1.1. Reliability statistics- Perceived Price scale (JASP output, generated by the author)

2. Perceived Agency (PA)

Frequentist Scale Reliability Statistics

Coefficient	Estimate	Std. Error	95% CI	
			Lower	Upper
McDonald's ω		0.039		
Cronbach's α	0.804	0.039	0.728	0.881

Note. Omega calculation with CFA failed. Try changing to PFA in 'Advanced Options'.

Frequentist Individual Item Reliability Statistics

Item	McDonald's ω (if item dropped)			Cronbach's α (if item dropped)		
	Estimate	Lower 95% CI	Upper 95% CI	Estimate	Lower 95% CI	Upper 95% CI
PA1_r				0.770	0.668	0.871
PA2_r				0.669	0.523	0.815
PA3_r				0.753	0.637	0.870

Note. Omega item dropped statistics with CFA failed.

Figure C1.2. Reliability statistics- Perceived Agency scale (JASP output, generated by the author)

3. Perceived Ethical Alignment (EA)

Frequentist Scale Reliability Statistics

Coefficient	Estimate	Std. Error	95% CI	
			Lower	Upper
McDonald's ω		0.026		
Cronbach's α	0.915	0.026	0.863	0.967

Note. Omega calculation with CFA failed. Try changing to PFA in 'Advanced Options'.

Frequentist Individual Item Reliability Statistics

Item	McDonald's ω (if item dropped)			Cronbach's α (if item dropped)		
	Estimate	Lower 95% CI	Upper 95% CI	Estimate	Lower 95% CI	Upper 95% CI
EA1_r				0.877	0.726	1.027
EA2_r				0.834	0.758	0.910
EA3_r				0.918	0.825	1.011

Note. Omega item dropped statistics with CFA failed.

Figure C1.3. Reliability statistics- Perceived Ethical Alignment scale (JASP output, generated by the author)

4. Perceived Platform Trust (PT)

Frequentist Scale Reliability Statistics

Coefficient	Estimate	Std. Error	95% CI	
			Lower	Upper
McDonald's ω		0.066		
Cronbach's α	0.783	0.066	0.653	0.912

Note. Omega calculation with CFA failed. Try changing to PFA in 'Advanced Options'.

Frequentist Individual Item Reliability Statistics

Item	McDonald's ω (if item dropped)			Cronbach's α (if item dropped)		
	Estimate	Lower 95% CI	Upper 95% CI	Estimate	Lower 95% CI	Upper 95% CI
PT1_r				0.798	0.580	1.016
PT2_r				0.508	0.213	0.803
PT3_r				0.775	0.678	0.872

Note. Omega item dropped statistics with CFA failed.

Figure C1.4. Reliability statistics- Perceived Platform Trust scale (JASP output, generated by the author)

5. Satisfaction (S)

Frequentist Scale Reliability Statistics

Coefficient	Estimate	Std. Error	95% CI	
			Lower	Upper
McDonald's ω		0.020		
Cronbach's α	0.894	0.020	0.855	0.933

Note. Omega calculation with CFA failed. Try changing to PFA in 'Advanced Options'.

Frequentist Individual Item Reliability Statistics

Item	McDonald's ω (if item dropped)			Cronbach's α (if item dropped)		
	Estimate	Lower 95% CI	Upper 95% CI	Estimate	Lower 95% CI	Upper 95% CI
S1_r				0.824	0.750	0.897
S2_r				0.806	0.728	0.885
S3_r				0.910	0.870	0.950

Note. Omega item dropped statistics with CFA failed.

Figure C1.5. Reliability statistics- Perceived Satisfaction scale (JASP output, generated by the author)

6. Algorithmic Literacy (AL)

Frequentist Scale Reliability Statistics

Coefficient	Estimate	Std. Error	95% CI	
			Lower	Upper
McDonald's ω		0.036		
Cronbach's α	0.801	0.036	0.730	0.871

Note. Omega calculation with CFA failed. Try changing to PFA in 'Advanced Options'.

Frequentist Individual Item Reliability Statistics

Item	McDonald's ω (if item dropped)			Cronbach's α (if item dropped)		
	Estimate	Lower 95% CI	Upper 95% CI	Estimate	Lower 95% CI	Upper 95% CI
AL2_r				0.732	0.629	0.835
AL1_r				0.772	0.672	0.872
AL3_r				0.673	0.543	0.802

Note. Omega item dropped statistics with CFA failed.

Figure C1.6. Reliability statistics- Algorithmic Literacy scale (JASP output, generated by the author)

7. Eco-ethical Values (EV)

Frequentist Scale Reliability Statistics

Coefficient	Estimate	Std. Error	95% CI	
			Lower	Upper
McDonald's ω		0.043		
Cronbach's α	0.743	0.043	0.658	0.828

Note. Omega calculation with CFA failed. Try changing to PFA in 'Advanced Options'.

Frequentist Individual Item Reliability Statistics

Item	McDonald's ω (if item dropped)			Cronbach's α (if item dropped)		
	Estimate	Lower 95% CI	Upper 95% CI	Estimate	Lower 95% CI	Upper 95% CI
EV1_r				0.515	0.330	0.699
EV2_r				0.541	0.369	0.714
EV3_r				0.833	0.761	0.905

Note. Omega item dropped statistics with CFA failed.

Figure C1.7. Reliability statistics- Eco-Ethical Values scale (JASP output, generated by the author)

C2 Exploratory Factor Analysis factor loading table

1. EFA factor loadings table

Kaiser-Meyer-Olkin Test

MSA	
Overall MSA	0.868
S1_r	0.883
S2_r	0.873
S3_r	0.944
AL1_r	0.760
AL2_r	0.835
AL3_r	0.668
PA1_r	0.806
PA2_r	0.814
PA3_r	0.922
EA1_r	0.877
EA2_r	0.876
EA3_r	0.927
PT1_r	0.872
PT2_r	0.842
PT3_r	0.840

Bartlett's Test

X ²	df	p
1188.046	105.000	< .001

Chi-Squared Test

	Value	df	p
Model	197.710	63	< .001

Factor Loadings

	Factor 1	Factor 2	Factor 3	Uniqueness
PT2_r	0.869			0.359
S1_r	0.851			0.319
S2_r	0.846			0.290
S3_r	0.700			0.375
PT1_r	0.690			0.548
EA1_r	0.677			0.352
EA2_r	0.661			0.307
EA3_r	0.651			0.338
PT3_r	0.589			0.608
PA2_r		0.820		0.335
PA1_r		0.707		0.544
PA3_r		0.632		0.316
AL3_r			0.885	0.246
AL1_r			0.681	0.546
AL2_r			0.679	0.385

Note. Applied rotation method is oblimin.

Factor Characteristics

	Unrotated solution			Rotated solution		
	Eigenvalue	Proportion var.	Cumulative	SumSq. Loadings	Proportion var.	Cumulative
Factor 1	6.550	0.437	0.437	5.288	0.353	0.353
Factor 2	1.536	0.102	0.539	2.004	0.134	0.486
Factor 3	1.045	0.070	0.609	1.839	0.123	0.609

Figure C2.1. EFA factor loadings table — 15-item three-factor solution, oblimin rotation (JASP output, generated by the author).

C3 Single-Factor Confirmatory Factor Analyses

1. Single-factor CFA — Perceived Agency

Model fit

Chi-square test

Model	X ²	df	p
Baseline model	459.686	3	
Factor model	0.000	0	

Note. The standard error method is robust.sem.

Additional fit measures

Fit indices

Index	Value
Comparative Fit Index (CFI)	1.000
Tucker-Lewis Index (TLI)	1.000
Bentler-Bonett Non-normed Fit Index (NNFI)	1.000
Bentler-Bonett Normed Fit Index (NFI)	1.000
Parsimony Normed Fit Index (PNFI)	0.000
Bollen's Relative Fit Index (RFI)	1.000
Bollen's Incremental Fit Index (IFI)	1.000
Relative Noncentrality Index (RNI)	1.000

Note. Except for the PNFI, the fit indices are scaled because of categorical variables in the data.

Other fit measures

Metric	Value
Root mean square error of approximation (RMSEA)	0.000
RMSEA 90% CI lower bound	0.000
RMSEA 90% CI upper bound	0.000
RMSEA p-value	
Standardized root mean square residual (SRMR)	1.086×10 ⁻⁹
Hoelter's critical N (α = .05)	
Hoelter's critical N (α = .01)	
Goodness of fit index (GFI)	1.000
McDonald fit index (MFI)	1.000
Expected cross validation index (ECVI)	

Note. The RMSEA results are scaled because of categorical variables in the data.

Figure C3.1. Single-factor CFA — Perceived Agency, WLSMV estimator (JASP output, generated by the author).

2. Single-factor CFA — Ethical Alignment

Model fit

Chi-square test

Model	X ²	df	p
Baseline model	1501.234	3	
Factor model	0.000	0	

Note. The standard error method is robust sem.

Additional fit measures

Fit indices

Index	Value
Comparative Fit Index (CFI)	1.000
Tucker-Lewis Index (TLI)	1.000
Bentler-Bonett Non-normed Fit Index (NNFI)	1.000
Bentler-Bonett Normed Fit Index (NFI)	1.000
Parsimony Normed Fit Index (PNFI)	0.000
Bollen's Relative Fit Index (RFI)	1.000
Bollen's Incremental Fit Index (IFI)	1.000
Relative Noncentrality Index (RNI)	1.000

Note. Except for the PNFI, the fit indices are scaled because of categorical variables in the data.

Other fit measures

Metric	Value
Root mean square error of approximation (RMSEA)	0.000
RMSEA 90% CI lower bound	0.000
RMSEA 90% CI upper bound	0.000
RMSEA p-value	
Standardized root mean square residual (SRMR)	6.933×10^{-10}
Hoelter's critical N ($\alpha = .05$)	
Hoelter's critical N ($\alpha = .01$)	
Goodness of fit index (GFI)	1.000
McDonald fit index (MFI)	1.000
Expected cross validation index (ECVI)	

Note. The RMSEA results are scaled because of categorical variables in the data.

Figure C3.2. Single-factor CFA — Ethical Alignment, WLSMV estimator (JASP output, generated by the author).

3. Single-factor CFA — Satisfaction

Model fit

Chi-square test

Model	χ^2	df	p
Baseline model	1923.638	3	
Factor model	0.000	0	

Note. The standard error method is robust.sem.

Additional fit measures

Fit indices

Index	Value
Comparative Fit Index (CFI)	1.000
Tucker-Lewis Index (TLI)	1.000
Bentler-Bonett Non-normed Fit Index (NNFI)	1.000
Bentler-Bonett Normed Fit Index (NFI)	1.000
Parsimony Normed Fit Index (PNFI)	0.000
Bollen's Relative Fit Index (RFI)	1.000
Bollen's Incremental Fit Index (IFI)	1.000
Relative Noncentrality Index (RNI)	1.000

Note. Except for the PNFI, the fit indices are scaled because of categorical variables in the data.

Other fit measures

Metric	Value
Root mean square error of approximation (RMSEA)	0.000
RMSEA 90% CI lower bound	0.000
RMSEA 90% CI upper bound	0.000
RMSEA p-value	
Standardized root mean square residual (SRMR)	2.739×10^{-9}
Hoelter's critical N ($\alpha = .05$)	
Hoelter's critical N ($\alpha = .01$)	
Goodness of fit index (GFI)	1.000
McDonald fit index (MFI)	1.000
Expected cross validation index (ECVI)	

Note. The RMSEA results are scaled because of categorical variables in the data.

Figure C3.3. Single-factor CFA — Satisfaction, WLSMV estimator (JASP output, generated by the author).

4. Single-factor CFA — Algorithmic Literacy

Model fit

Chi-square test

Model	X ²	df	p
Baseline model	383.050	3	
Factor model	0.000	0	

Note. The standard error method is robust.sem.

Additional fit measures

Fit indices

Index	Value
Comparative Fit Index (CFI)	1.000
Tucker-Lewis Index (TLI)	1.000
Bentler-Bonett Non-normed Fit Index (NNFI)	1.000
Bentler-Bonett Normed Fit Index (NFI)	1.000
Parsimony Normed Fit Index (PNFI)	0.000
Bollen's Relative Fit Index (RFI)	1.000
Bollen's Incremental Fit Index (IFI)	1.000
Relative Noncentrality Index (RNI)	1.000

Note. Except for the PNFI, the fit indices are scaled because of categorical variables in the data.

Other fit measures

Metric	Value
Root mean square error of approximation (RMSEA)	0.000
RMSEA 90% CI lower bound	0.000
RMSEA 90% CI upper bound	0.000
RMSEA p-value	
Standardized root mean square residual (SRMR)	1.595×10 ⁻⁹
Hoelter's critical N (α = .05)	
Hoelter's critical N (α = .01)	
Goodness of fit index (GFI)	1.000
McDonald fit index (MFI)	1.000
Expected cross validation index (ECVI)	

Note. The RMSEA results are scaled because of categorical variables in the data.

Figure C3.4. Single-factor CFA — Algorithmic Literacy, WLSMV estimator (JASP output, generated by the author).

5. Single-factor CFA — Platform Trust

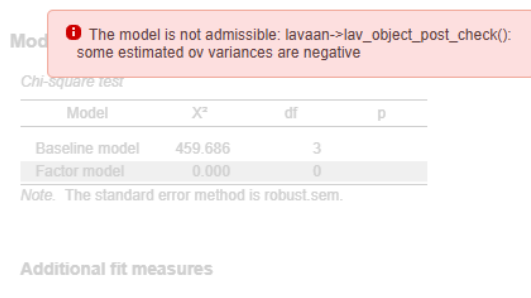


Figure C3.5. Single-factor CFA — Platform Trust, WLSMV estimator (JASP output, generated by the author).

C4 Full descriptive statistics output, including distribution plots

1. Descriptive statistics table (PA, EA, PT, S, AL, EV, composites)

Descriptive Statistics

	PA_comp	EA_comp	PT_comp	S_comp	AL_comp	EV_comp
Valid	122	122	122	122	122	122
Missing	0	0	0	0	0	0
Mean (arithmetic)	4.391	4.495	3.948	4.448	4.005	4.995
Std. Deviation	1.380	1.376	1.261	1.312	1.455	1.187
Skewness	-0.520	-1.073	-0.421	-0.824	-0.157	-0.752
Std. Error of Skewness	0.219	0.219	0.219	0.219	0.219	0.219
Kurtosis	-0.532	0.286	-0.430	-0.059	-0.761	0.659
Std. Error of Kurtosis	0.435	0.435	0.435	0.435	0.435	0.435
Shapiro-Wilk	0.949	0.869	0.968	0.923	0.977	0.952
P-value of Shapiro-Wilk	< .001	< .001	.005	< .001	.033	< .001
Minimum	1.000	1.000	1.000	1.000	1.000	1.000
Maximum	7.000	6.000	6.333	6.667	7.000	7.000

Figure C4.1. Descriptive statistics- PA, EA, PT, S, AL, EV composites (JASP output, generated by the author)

2. Distribution plots (PA, EA, PT, S, AL, EV)

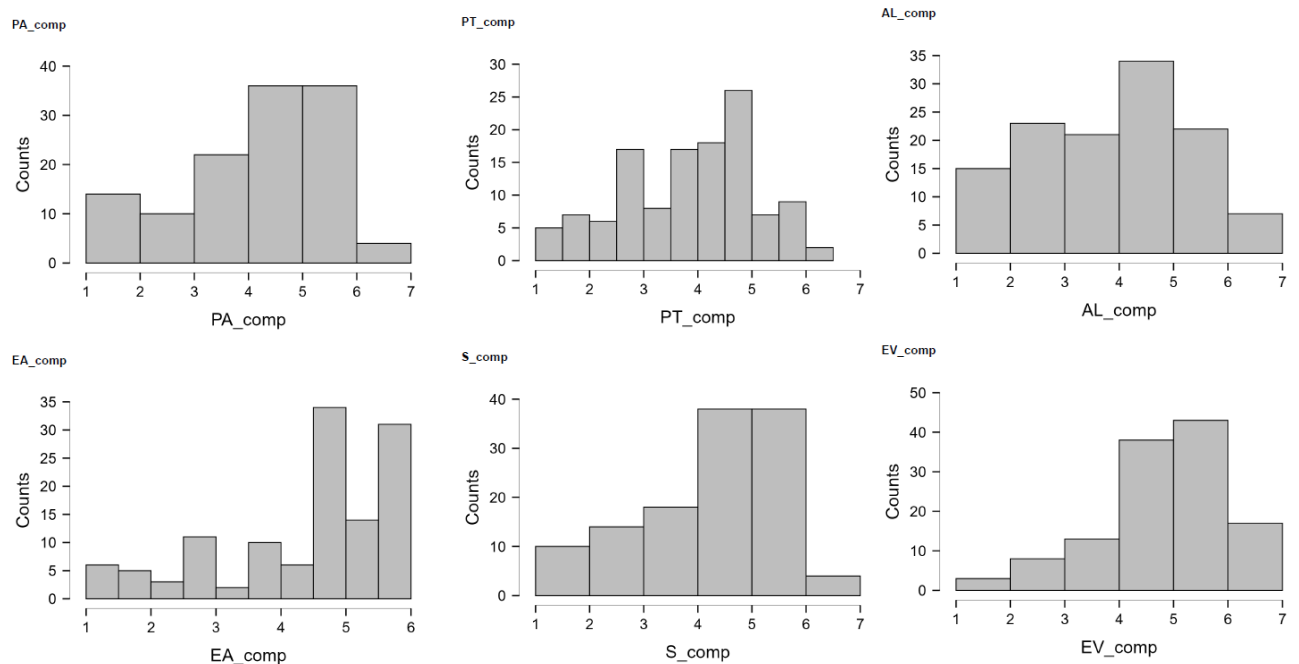


Figure C4.2. Distribution plots- PA, EA, PT, S, AL, EV composites (JASP output, generated by the author)

C5 Frequency tables for all constructs

1. Frequency table- Age (Control and Treatment)

Group	Age	Frequency	Percent	Valid Percent	Cumulative Percent
Control	18-24	10	17.2	17.2	17.2
	25-34	11	19.0	19.0	36.2
	35-44	9	15.5	15.5	51.7
	45-54	20	34.5	34.5	86.2
	55-64	5	8.6	8.6	94.8
	65 and older	3	5.2	5.2	100.0
	Missing	0	0.0		
	Total	58	100.0		
Treatment	18-24	11	17.2	17.5	17.5
	25-34	11	17.2	17.5	34.9
	35-44	10	15.6	15.9	50.8
	45-54	11	17.2	17.5	68.3
	55-64	17	26.6	27.0	95.2
	65 and older	3	4.7	4.8	100.0
	Missing	1	1.6		
	Total	64	100.0		

Figure C5.1. Frequency table- Age (JASP output, generated by the author)

2. Frequency table- Gender (Control and Treatment)

Group	Gender	Frequency	Percent	Valid Percent	Cumulative Percent
Control	Female	36	62.1	62.1	62.1
	Male	22	37.9	37.9	100.0
	Prefer not to say	0	0.0	0.0	100.0
	Missing	0	0.0		
	Total	58	100.0		
Treatment	Female	47	73.4	74.6	74.6
	Male	15	23.4	23.8	98.4
	Prefer not to say	1	1.6	1.6	100.0
	Missing	1	1.6		
	Total	64	100.0		

Figure C5.2. Frequency table- Gender (JASP output, generated by the author)

3. Frequency table- Country (Control and Treatment)

Group	Country	Frequency	Percent	Valid Percent	Cumulative Percent
Control	Belgium	49	84.5	84.5	84.5
	Germany	2	3.4	3.4	87.9
	Luxembourg	5	8.6	8.6	96.6
	Other	1	1.7	1.7	98.3
	U.S.	1	1.7	1.7	100.0
	Missing	0	0.0		
	Total	58	100.0		
Treatment	Belgium	53	82.8	84.1	84.1
	Germany	6	9.4	9.5	93.7
	Luxembourg	4	6.3	6.3	100.0
	Other	0	0.0	0.0	100.0
	U.S.	0	0.0	0.0	100.0
	Missing	1	1.6		
	Total	64	100.0		

Figure C5.3. Frequency table- Country (JASP output, generated by the author)

4. Frequency table- Employment (Control and Treatment)

Group	Employment	Frequency	Percent	Valid Percent	Cumulative Percent
Control	Employed full-time	29	50.0	50.0	50.0
	Employed part-time	16	27.6	27.6	77.6
	Prefer not to say	0	0.0	0.0	77.6
	Retired	3	5.2	5.2	82.8
	Self-employed / Freelancer	2	3.4	3.4	86.2
	Student	8	13.8	13.8	100.0
	Unable to work	0	0.0	0.0	100.0
	Unemployed / Looking for work	0	0.0	0.0	100.0
	Missing	0	0.0		
	Total	58	100.0		
Treatment	Employed full-time	23	35.9	36.5	36.5
	Employed part-time	18	28.1	28.6	65.1
	Prefer not to say	2	3.1	3.2	68.3
	Retired	3	4.7	4.8	73.0
	Self-employed / Freelancer	6	9.4	9.5	82.5
	Student	9	14.1	14.3	96.8
	Unable to work	1	1.6	1.6	98.4
	Unemployed / Looking for work	1	1.6	1.6	100.0
	Missing	1	1.6		
	Total	64	100.0		

Figure C5.4. Frequency table- Employment (JASP output, generated by the author)

5. Frequency table- Experience (Control and Treatment)

Group	Experience	Frequency	Percent	Valid Percent	Cumulative Percent
Control	3 to 5 times	17	29.3	29.3	29.3
	More than 5 times	14	24.1	24.1	53.4
	Never	5	8.6	8.6	62.1
	One to two times	22	37.9	37.9	100.0
	Missing	0	0.0		
	Total	58	100.0		
Treatment	3 to 5 times	16	25.0	25.0	25.0
	More than 5 times	13	20.3	20.3	45.3
	Never	6	9.4	9.4	54.7
	One to two times	29	45.3	45.3	100.0
	Missing	0	0.0		
	Total	64	100.0		

Figure C5.5. Frequency table- Experience (JASP output, generated by the author)

C6 Correlations

1. Full Pearson correlation matrix

Pearson's Correlations

			Pearson's r	p	Lower 95% CI	Upper 95% CI
PA_comp	-	EA_comp	0.558***	< .001	0.422	0.669
PA_comp	-	PT_comp	0.425***	< .001	0.267	0.560
PA_comp	-	EV_comp	-0.063	.490	-0.238	0.116
PA_comp	-	S_comp	0.463***	< .001	0.310	0.592
PA_comp	-	BI1_r	0.329***	< .001	0.161	0.479
PA_comp	-	AL_comp	0.190*	.036	0.013	0.356
PA_comp	-	PS1_r	-0.096	.297	-0.269	0.084
PA_comp	-	PS2_r	0.019	.838	-0.160	0.197
EA_comp	-	PT_comp	0.652***	< .001	0.536	0.743
EA_comp	-	EV_comp	-0.219*	.015	-0.382	-0.043
EA_comp	-	S_comp	0.721***	< .001	0.623	0.797
EA_comp	-	BI1_r	0.553***	< .001	0.416	0.665
EA_comp	-	AL_comp	0.292**	.001	0.120	0.446
EA_comp	-	PS1_r	-0.261**	.004	-0.420	-0.086
EA_comp	-	PS2_r	-0.045	.626	-0.221	0.135
PT_comp	-	EV_comp	-0.127	.163	-0.298	0.052
PT_comp	-	S_comp	0.710***	< .001	0.610	0.789
PT_comp	-	BI1_r	0.537***	< .001	0.397	0.653
PT_comp	-	AL_comp	0.221*	.015	0.045	0.384
PT_comp	-	PS1_r	-0.227*	.012	-0.390	-0.051
PT_comp	-	PS2_r	-0.074	.422	-0.249	0.106
EV_comp	-	S_comp	-0.219*	.015	-0.382	-0.043
EV_comp	-	BI1_r	-0.302***	< .001	-0.456	-0.132
EV_comp	-	AL_comp	-0.271**	.003	-0.428	-0.098
EV_comp	-	PS1_r	0.492***	< .001	0.343	0.616
EV_comp	-	PS2_r	0.014	.880	-0.165	0.192
S_comp	-	BI1_r	0.697***	< .001	0.593	0.779
S_comp	-	AL_comp	0.328***	< .001	0.159	0.478
S_comp	-	PS1_r	-0.215*	.018	-0.379	-0.038
S_comp	-	PS2_r	-0.025	.785	-0.203	0.154
BI1_r	-	AL_comp	0.345***	< .001	0.178	0.492
BI1_r	-	PS1_r	-0.252**	.005	-0.412	-0.077
BI1_r	-	PS2_r	0.010	.917	-0.169	0.188
AL_comp	-	PS1_r	-0.025	.781	-0.203	0.154
AL_comp	-	PS2_r	0.147	.107	-0.032	0.317
PS1_r	-	PS2_r	0.210*	.021	0.033	0.374

* p < .05, ** p < .01, *** p < .001

Figure C6.1. Correlation Matrix (JASP output, generated by the author)

C7 Assumption Checks

Test of Equality of Variances (Levene's)

	F	df ₁	df ₂	p
PA_comp	0.541	1	120	.464
EA_comp	3.731	1	120	.056
PT_comp	0.306	1	120	.581
S_comp	1.534	1	120	.218
AL_comp	0.136	1	120	.713

Figure C7.1. Assumption Checks (JASP output, generated by the author)

C8 Hypothesis Testing

1. T-test: Perceived Agency- H2

Independent Samples T-Test

	Test	Statistic	df	p	Cohen's d	SE Cohen's d
PA_comp	Student	-3.176	120.0	.002	-0.576	0.189
	Welch	-3.160	115.3	.002	-0.574	0.189

Descriptives

Group Descriptives

	Group	N	Mean	SD	SE	Coefficient of variation
PA_comp	Control	58	3.989	1.402	0.184	0.351
	Treatment	64	4.755	1.264	0.158	0.266

Figure C8.1. T-test: Perceived Agency (JASP output, generated by the author)

2. T-test: Ethical Alignment- H4

Independent Samples T-Test

	Test	Statistic	df	p	Cohen's d	SE Cohen's d
EA_comp	Student	-1.013	120.0	.313	-0.184	0.182
	Welch	-1.005	112.8	.317	-0.183	0.182

Descriptives

Group Descriptives

	Group	N	Mean	SD	SE	Coefficient of variation
EA_comp	Control	58	4.362	1.483	0.195	0.340
	Treatment	64	4.615	1.270	0.159	0.275

Figure C8.2. T-test: Ethical Alignment (JASP output, generated by the author)

3. T-test: Platform Trust- H3

Independent Samples T-Test

	Test	Statistic	df	p	Cohen's d	SE Cohen's d
PT_comp	Student	1.153	120.0	.251	0.209	0.182
	Welch	1.150	117.1	.253	0.209	0.182

Descriptives

Group Descriptives

	Group	N	Mean	SD	SE	Coefficient of variation
PT_comp	Control	58	4.086	1.297	0.170	0.317
	Treatment	64	3.823	1.224	0.153	0.320

Figure C8.3. T-test: Platform Trust (JASP output, generated by the author)

4. T-test: Satisfaction- H1

Independent Samples T-Test

	Test	Statistic	df	p	Cohen's d	SE Cohen's d
S_comp	Student	0.645	120.0	.520	0.117	0.182
	Welch	0.640	113.2	.524	0.116	0.182

Descriptives

Group Descriptives

	Group	N	Mean	SD	SE	Coefficient of variation
S_comp	Control	58	4.529	1.414	0.186	0.312
	Treatment	64	4.375	1.219	0.152	0.279

Figure C8.4. T-test: Satisfaction (JASP output, generated by the author)

5. T-test: Algorithmic Literacy- covariate check

Independent Samples T-Test

	Test	Statistic	df	p	Cohen's d	SE Cohen's d
AL_comp	Student	2.108	120.0	.037	0.382	0.185
	Welch	2.107	118.7	.037	0.382	0.185

Descriptives

Group Descriptives

	Group	N	Mean	SD	SE	Coefficient of variation
AL_comp	Control	58	4.293	1.439	0.189	0.335
	Treatment	64	3.745	1.431	0.179	0.382

Figure C8.5. T-test: Algorithmic Literacy (JASP output, generated by the author)

6. Hierarchical regression predicting Satisfaction- H1 full model

Model Summary - S_comp

Model	R	R ²	Adjusted R ²	RMSE	R ² Change	df1	df2	p
M ₀	0.045	0.002	-0.006	1.303	0.002	2	119	.887
M ₁	0.390	0.152	0.107	1.227	0.150	6	114	.004

Note. M₀ includes Group_d:Group_d

Note. M₁ includes Group_d:Group_d, Include:Include, AL_comp, PS1_r, PS2_r, EV_comp

ANOVA

Model		Sum of Squares	df	Mean Square	F	p
M ₀	Regression	0.407	1	0.407	0.240	.625
	Residual	201.9	119	1.697		
	Total	202.4	120			
M ₁	Regression	30.72	6	5.120	3.400	.004
	Residual	171.6	114	1.506		
	Total	202.4	120			

Note. M₀ includes Group_d:Group_d

Note. M₁ includes Group_d:Group_d, Include:Include, AL_comp, PS1_r, PS2_r, EV_comp

Coefficients

Model		Unstandardized	Standard Error	Standardized*	t	p
M ₀	(Intercept)	4.529	0.171		26.48	< .001
	Group_d (0) *					
	Group_d (1) *	NaN	NaN		NaN	NaN
	Group_d (1) *	-0.116	0.237		-0.489	.625
M ₁	(Intercept)	4.584	0.835		5.490	< .001
	AL_comp	0.262	0.085	0.294	3.090	.003
	PS1_r	-0.121	0.083	-0.152	-1.465	.146
	PS2_r	-0.038	0.100	-0.036	-0.383	.703
	EV_comp	-0.107	0.117	-0.096	-0.909	.366
	Group_d (0) *					
	Group_d (1) *	NaN	NaN		NaN	NaN
	Group_d (1) *	0.009	0.236		0.036	.971
	Include (0) *					
	Include (1) *	NaN	NaN		NaN	NaN
	Include (1) *	0.013	0.257		0.050	.960

Note. Missing coefficients are undefined because of singularities. Check the data for anything out of order!

* Standardized coefficients can only be computed for continuous predictors.

Figure C8.6. Hierarchical regression predicting Satisfaction- H1 full model (JASP output, generated by the author)

7. Mediation output: PA as a mediator- H2

Parameter estimates

Direct effects

					95% Confidence Interval	
	Estimate	Std. error	z-value	p	Lower	Upper
Group_d → S_comp	-0.505	0.212	-2.383	.017	-0.912	-0.088

Note. Estimator is ML.

Indirect effects

					95% Confidence Interval	
	Estimate	Std. error	z-value	p	Lower	Upper
Group_d → PA_comp → S_comp	0.364	0.124	2.932	.003	0.135	0.680

Note. Estimator is ML.

Total effects

					95% Confidence Interval	
	Estimate	Std. error	z-value	p	Lower	Upper
Group_d → S_comp	-0.141	0.228	-0.624	.533	-0.618	0.307

Note. Estimator is ML.

Path coefficients

					95% Confidence Interval	
	Estimate	Std. error	z-value	p	Lower	Upper
PA_comp → S_comp	0.457	0.088	5.172	< .001	0.273	0.619
Group_d → S_comp	-0.505	0.212	-2.383	.017	-0.912	-0.088
Group_d → PA_comp	0.796	0.244	3.263	.001	0.316	1.276
PS1_r → Group_d	6.835×10^{-4}	0.034	0.020	.984	-0.067	0.068
PS2_r → Group_d	0.014	0.039	0.370	.711	-0.066	0.086
EV_comp → Group_d	-0.018	0.048	-0.380	.704	-0.119	0.070
PS1_r → PA_comp	-0.066	0.085	-0.768	.442	-0.231	0.103
PS2_r → PA_comp	0.028	0.102	0.278	.781	-0.171	0.232
EV_comp → PA_comp	-0.047	0.123	-0.381	.704	-0.304	0.177
PS1_r → S_comp	-0.063	0.077	-0.822	.411	-0.219	0.086
PS2_r → S_comp	-0.008	0.086	-0.098	.922	-0.171	0.170
EV_comp → S_comp	-0.202	0.101	-1.998	.046	-0.385	0.018

Note. Estimator is ML.

Figure C8.7. Mediation output: PA as a mediator- H2 (JASP output, generated by the author)

8. Mediation output: PT as a mediator- H3

Parameter estimates

Direct effects

					95% Confidence Interval	
					Lower	Upper
Group_d	→	S_comp	Estimate	Std. error	z-value	p
			0.036	0.173	0.205	.837
						-0.294
						0.397

Note. Estimator is ML.

Indirect effects

						95% Confidence Interval				
						Lower	Upper			
Group_d	→	PT_comp	→	S_comp	Estimate	Std. error	z-value	p		
					-0.177	0.158	-1.120	.263		
									-0.508	0.145

Note. Estimator is ML.

Total effects

					95% Confidence Interval		
					Lower	Upper	
Group_d	→	S_comp	Estimate	Std. error	z-value	p	
			-0.141	0.226	-0.624	.533	-0.594 0.327

Note. Estimator is ML.

Path coefficients

						95% Confidence Interval		
						Lower	Upper	
		Estimate	Std. error	z-value	p			
PT_comp	→	S_comp	0.710	0.067	10.66	< .001	0.566	0.828
Group_d	→	S_comp	0.036	0.173	0.205	.837	-0.294	0.397
Group_d	→	PT_comp	-0.249	0.229	-1.084	.278	-0.691	0.206
PS1_r	→	Group_d	8.835×10^{-4}	0.034	0.020	.984	-0.068	0.067
PS2_r	→	Group_d	0.014	0.038	0.376	.707	-0.063	0.087
EV_comp	→	Group_d	-0.018	0.048	-0.379	.705	-0.119	0.070
PS1_r	→	PT_comp	-0.148	0.084	-1.773	.076	-0.312	0.016
PS2_r	→	PT_comp	-0.029	0.077	-0.382	.702	-0.181	0.122
EV_comp	→	PT_comp	-0.065	0.118	-0.550	.582	-0.307	0.153
PS1_r	→	S_comp	0.012	0.060	0.202	.840	-0.108	0.126
PS2_r	→	S_comp	0.025	0.062	0.408	.683	-0.096	0.153
EV_comp	→	S_comp	-0.177	0.098	-1.805	.071	-0.354	0.027

Note. Estimator is ML.

Figure C8.8. Mediation output: PT as a mediator- H3 (JASP output, generated by the author)

9. Mediation output: EA as a mediator- H4

Parameter estimates

Direct effects

		Estimate	Std. error	z-value	p	95% Confidence Interval	
						Lower	Upper
Group_d	→ S_comp	-0.320	0.168	-1.901	.057	-0.660	8.868×10 ⁻⁵

Note. Estimator is ML.

Indirect effects

		Estimate	Std. error	z-value	p	95% Confidence Interval	
						Lower	Upper
Group_d	→ EA_comp → S_comp	0.179	0.159	1.124	.261	-0.118	0.517

Note. Estimator is ML.

Total effects

		Estimate	Std. error	z-value	p	95% Confidence Interval	
						Lower	Upper
Group_d	→ S_comp	-0.141	0.228	-0.624	.533	-0.592	0.309

Note. Estimator is ML.

Path coefficients

		Estimate	Std. error	z-value	p	95% Confidence Interval	
						Lower	Upper
EA_comp	→ S_comp	0.673	0.066	10.23	< .001	0.544	0.800
Group_d	→ S_comp	-0.320	0.168	-1.901	.057	-0.660	8.868×10 ⁻⁵
Group_d	→ EA_comp	0.266	0.240	1.110	.267	-0.181	0.749
PS1_r	→ Group_d	6.835×10 ⁻⁴	0.035	0.020	.984	-0.069	0.069
PS2_r	→ Group_d	0.014	0.039	0.372	.710	-0.065	0.087
EV_comp	→ Group_d	-0.018	0.048	-0.377	.706	-0.115	0.073
PS1_r	→ EA_comp	-0.150	0.091	-1.657	.097	-0.321	0.035
PS2_r	→ EA_comp	-0.009	0.086	-0.108	.914	-0.171	0.165
EV_comp	→ EA_comp	-0.188	0.149	-1.250	.211	-0.497	0.079
PS1_r	→ S_comp	0.008	0.065	0.117	.907	-0.118	0.141
PS2_r	→ S_comp	0.011	0.061	0.178	.859	-0.109	0.129
EV_comp	→ S_comp	-0.098	0.110	-0.896	.370	-0.297	0.133

Note. Estimator is ML.

Figure C8.9. Mediation output: EA as a mediator- H4 (JASP output, generated by the author)

10. Moderation output (PROCESS Model 1): Group x AL- H5

Path coefficients

							95% Confidence Interval	
			Estimate	Std. Error	z-value	p	Lower	Upper
Group_d	→	S_comp	0.007	0.228	0.029	.977	-0.440	0.453
AL_comp	→	S_comp	0.296	0.078	3.773	< .001	0.142	0.450
Group_d:AL_comp	→	S_comp	0.073	0.157	0.462	.644	-0.235	0.381

Note. Moderation effect estimates are based on mean-centered variables.

Direct and indirect effects

							95% Confidence Interval		
							Lower	Upper	
Group_d	→	S_comp	16	-0.115	0.351	-0.327	.744	-0.802	0.572
Group_d	→	S_comp	50	0.031	0.233	0.131	.896	-0.425	0.487
Group_d	→	S_comp	84	0.103	0.306	0.337	.736	-0.497	0.704

Note. Moderation effect estimates are based on mean-centered variables.

Total effects

									95% Confidence Interval		
					AL_comp	Estimate	Std. Error	z-value	p	Lower	Upper
Total	Group_d	→	S_comp	16		-0.115	0.351	-0.327	.744	-0.802	0.572
	Group_d	→	S_comp	50		0.031	0.233	0.131	.896	-0.425	0.487
	Group_d	→	S_comp	84		0.103	0.306	0.337	.736	-0.497	0.704

Note. Moderation effect estimates are based on mean-centered variables.

Figure C8.10. Moderation output (PROCESS Model 1) (JASP output, generated by the author)

11. Moderation mediation output (PROCESS Model 7): Perceived Agency

Outcome Variable: PA_comp

Model Summary:

R	R-sq	MSE	F	df1	df2	p
0.4038	0.1631	1.6763	3.1452	7.0000	113.0000	0.0045

Model:

	coeff	se	t	p	LLCI	ULCI
constant	3.9584	0.7581	5.2212	0.0000	2.4564	5.4605
Group_d	0.9125	0.2495	3.6581	0.0004	0.4183	1.4067
AL_centred	0.3463	0.1258	2.7533	0.0069	0.0971	0.5955
int_1	-0.1878	0.1684	-1.1151	0.2672	-0.5214	0.1459
EV_comp	0.0584	0.1239	0.4715	0.6382	-0.1870	0.3038
Include	0.1305	0.2740	0.4762	0.6348	-0.4124	0.6734
PS1_r	-0.0943	0.0875	-1.0780	0.2833	-0.2676	0.0790
PS2_r	-0.0158	0.1061	-0.1485	0.8822	-0.2260	0.1945

Product terms key:

int_1 : Group_d x AL_centred

Test(s) of highest order unconditional interaction(s):

R2-chng	F	df1	df2	p
X*W	0.0092	1.2434	1.0000	113.0000

Outcome Variable: S_comp

Model Summary:

R	R-sq	MSE	F	df1	df2	p
0.5402	0.2918	1.2571	7.8277	6.0000	114.0000	0.0000

Model:

	coeff	se	t	p	LLCI	ULCI
constant	3.9745	0.7314	5.4344	0.0000	2.5257	5.4234
Group_d	-0.4809	0.2220	-2.1661	0.0324	-0.9206	-0.0411
PA_comp	0.4571	0.0784	5.8282	0.0000	0.3017	0.6125
EV_comp	-0.1953	0.1024	-1.9082	0.0589	-0.3981	0.0074
Include	-0.0931	0.2323	-0.4008	0.6893	-0.5533	0.3670
PS1_r	-0.0688	0.0755	-0.9122	0.3636	-0.2184	0.0807
PS2_r	0.0031	0.0903	0.0344	0.9726	-0.1758	0.1821

Direct effect of X on Y:

effect	se	t	p	LLCI	ULCI
-0.4809	0.2220	-2.1661	0.0324	-0.9206	-0.0411

Conditional indirect effects of X on Y:

INDIRECT EFFECT:

Group_d -> PA_comp -> S_comp

AL_centred	Effect	BootSE	BootLLCI	BootULCI
-1.6721	0.5607	0.2125	0.1746	1.0069
0.3279	0.3890	0.1517	0.1206	0.7128
1.4879	0.2894	0.1935	-0.0608	0.7005

Index of moderated mediation:

Index	BootSE	BootLLCI	BootULCI
AL_centred	-0.0858	0.0862	-0.2553

Figure C8.11. Exploratory moderated mediation output — Model A: Perceived Agency as mediator, Algorithmic Literacy as moderator (PROCESS Model 7, Hayes, 2022; R Version 5.0). Bootstrap resamples = 5,000, seed = 1234. N = 121.

12. Moderation mediation output (PROCESS Model 7): Perceived Trust

Outcome Variable: PT_comp

Model Summary:

R	R-sq	MSE	F	df1	df2	p
0.3363	0.1131	1.4898	2.0589	7.0000	113.0000	0.0537

Model:

	coeff	se	t	p	LLCI	ULCI
constant	4.7733	0.7147	6.6785	0.0000	3.3573	6.1894
Group_d	-0.0889	0.2352	-0.3782	0.7060	-0.5548	0.3770
AL_centred	0.2284	0.1186	1.9257	0.0567	-0.0066	0.4633
int_1	-0.1123	0.1588	-0.7076	0.4807	-0.4269	0.2022
EV_comp	0.0296	0.1168	0.2538	0.8001	-0.2017	0.2610
Include	-0.2232	0.2583	-0.8640	0.3894	-0.7351	0.2886
PS1_r	-0.1857	0.0825	-2.2520	0.0263	-0.3490	-0.0223
PS2_r	-0.0219	0.1001	-0.2191	0.8269	-0.2201	0.1763

Product terms key:

int_1 : Group_d x AL_centred

Test(s) of highest order unconditional interaction(s):

R2-chng	F	df1	df2	p	
X*W	0.0039	0.5007	1.0000	113.0000	0.4807

Outcome Variable: S_comp

Model Summary:

R	R-sq	MSE	F	df1	df2	p
0.7239	0.5240	0.8449	20.9181	6.0000	114.0000	0.0000

Model:

	coeff	se	t	p	LLCI	ULCI
constant	2.3598	0.6354	3.7138	0.0003	1.1010	3.6186
Group_d	0.0038	0.1749	0.0218	0.9826	-0.3427	0.3503
PT_comp	0.7157	0.0695	10.3038	0.0000	0.5781	0.8533
EV_comp	-0.1865	0.0839	-2.2228	0.0282	-0.3527	-0.0203
Include	0.1286	0.1918	0.6706	0.5038	-0.2513	0.5085
PS1_r	0.0205	0.0628	0.3263	0.7448	-0.1039	0.1449
PS2_r	0.0097	0.0740	0.1314	0.8957	-0.1369	0.1564

Direct effect of X on Y:

effect	se	t	p	LLCI	ULCI
0.0038	0.1749	0.0218	0.9826	-0.3427	0.3503

Conditional indirect effects of X on Y:

INDIRECT EFFECT:

Group_d -> PT_comp -> S_comp

AL_centred	Effect	BootSE	BootLLCI	BootULCI
-1.6721	0.0708	0.2668	-0.4941	0.5692
0.3279	-0.0900	0.1896	-0.4558	0.2960
1.4879	-0.1833	0.2650	-0.7005	0.3526

Index of moderated mediation:

Index	BootSE	BootLLCI	BootULCI	
AL_centred	-0.0804	0.1222	-0.3077	0.1753

Figure C8.12. Exploratory moderated mediation output — Model B: Platform Trust as mediator, Algorithmic Literacy as moderator (PROCESS Model 7, Hayes, 2022; R Version 5.0). Bootstrap resamples = 5,000, seed = 1234. N = 121.

13. Moderation mediation output (PROCESS Model 7): Ethical Alignment

Outcome Variable: EA_comp

Model Summary:

	R	R-sq	MSE	F	df1	df2	p
	0.4226	0.1786	1.6300	3.5096	7.0000	113.0000	0.0019

Model:

	coeff	se	t	p	LLCI	ULCI
constant	5.4532	0.7476	7.2943	0.0000	3.9721	6.9343
Group_d	0.4777	0.2460	1.9419	0.0546	-0.0097	0.9650
AL_centred	0.3029	0.1240	2.4423	0.0161	0.0572	0.5487
int_1	-0.0618	0.1661	-0.3720	0.7106	-0.3908	0.2672
EV_comp	-0.0485	0.1221	-0.3971	0.6921	-0.2905	0.1935
Include	-0.2001	0.2702	-0.7404	0.4606	-0.7354	0.3353
PS1_r	-0.1945	0.0862	-2.2556	0.0260	-0.3654	-0.0237
PS2_r	-0.0235	0.1047	-0.2250	0.8224	-0.2309	0.1838

Product terms key:
int_1 : Group_d x AL_centred

Test(s) of highest order unconditional interaction(s):

	R2-chng	F	df1	df2	p
X*W	0.0010	0.1384	1.0000	113.0000	0.7106

Outcome Variable: S_comp

Model Summary:

	R	R-sq	MSE	F	df1	df2	p
	0.7306	0.5338	0.8276	21.7512	6.0000	114.0000	0.0000

Model:

	coeff	se	t	p	LLCI	ULCI
constant	2.0202	0.6450	3.1324	0.0022	0.7426	3.2979
Group_d	-0.3532	0.1742	-2.0273	0.0450	-0.6983	-0.0081
EA_comp	0.6775	0.0644	10.5244	0.0000	0.5499	0.8050
EV_comp	-0.1066	0.0836	-1.2747	0.2050	-0.2723	0.0591
Include	0.1230	0.1897	0.6485	0.5180	-0.2528	0.4988
PS1_r	0.0157	0.0621	0.2524	0.8012	-0.1073	0.1386
PS2_r	-0.0044	0.0733	-0.0602	0.9521	-0.1496	0.1408

Direct effect of X on Y:

	effect	se	t	p	LLCI	ULCI
	-0.3532	0.1742	-2.0273	0.0450	-0.6983	-0.0081

Conditional indirect effects of X on Y:

INDIRECT EFFECT:

Group_d -> EA_comp -> S_comp

	AL_centred	Effect	BootSE	BootLLCI	BootULCI
	-1.6721	0.3936	0.2834	-0.1752	0.9383
	0.3279	0.3099	0.1945	-0.0561	0.7153
	1.4879	0.2613	0.2480	-0.2109	0.7628

Index of moderated mediation:

	Index	BootSE	BootLLCI	BootULCI
AL_centred	-0.0419	0.1148	-0.2638	0.1859

Figure C8.13. Exploratory moderated mediation output — Model C: Ethical Alignment as mediator, Algorithmic Literacy as moderator (PROCESS Model 7, Hayes, 2022; R Version 5.0). Bootstrap resamples = 5,000, seed = 1234. N = 121.

C9 Exploratory Analysis

1. T-test: Behavioural Intention (BI1_r)

Independent Samples T-Test

	Test	Statistic	df	p	Cohen's d	SE Cohen's d
BI1_r	Student	1.085	120.0	.280	0.197	0.182
	Welch	1.092	119.9	.277	0.197	0.182

Assumption Checks

Test of Normality (Shapiro-Wilk)

Residuals	W	p
BI1_r	0.914	< .001

Note. Significant results suggest a deviation from normality.

Test of Equality of Variances (Levene's)

	F	df ₁	df ₂	p
BI1_r	1.251	1	120	.266

Descriptives

Group Descriptives

	Group	N	Mean	SD	SE	Coefficient of variation
BI1_r	Control	58	4.810	1.468	0.193	0.305
	Treatment	64	4.500	1.671	0.209	0.371

Figure C9.1. T-test: Behavioural Intention (BI1_r) (JASP output, generated by the author)

List of Resource Persons

The following individuals provided guidance, feedback, or expertise that contributed to the development of this thesis:

El Midaoui Youssra – Thesis Supervisor, HEC Liège, University of Liège

Steils Nadia – Reading Committee Member, HEC Liège, University of Liège

References

- Afroogh, S., Akbari, A., Malone, E., Kargar, M., & Alambeigi, H. (2024). Trust in AI: Progress, challenges, and future directions. *Humanities and Social Sciences Communications*, 11(1), 1568. <https://doi.org/10.1057/s41599-024-04044-8>
- Aguinis, H., & Bradley, K. J. (2014). Best Practice Recommendations for Designing and Implementing Experimental Vignette Methodology Studies. *Organizational Research Methods*, 17(4), 351–371. <https://doi.org/10.1177/1094428114547952>
- Aguirre, E., Mahr, D., Grewal, D., De Ruyter, K., & Wetzels, M. (2015). Unraveling the Personalization Paradox: The Effect of Information Collection and Trust-Building Strategies on Online Advertisement Effectiveness. *Journal of Retailing*, 91(1), 34–49. <https://doi.org/10.1016/j.jretai.2014.09.005>
- Alderighi, M., Nava, C. R., Calabrese, M., Christille, J.-M., & Salvemini, C. B. (2022). Consumer perception of price fairness and dynamic pricing: Evidence from Booking.com. *Journal of Business Research*, 145, 769–783. <https://doi.org/10.1016/j.jbusres.2022.03.017>
- Ali, F., Ryu, K., & Hussain, K. (2016). Influence of Experiences on Memories, Satisfaction and Behavioral Intentions: A Study of Creative Tourism. *Journal of Travel & Tourism Marketing*, 33(1), 85–100. <https://doi.org/10.1080/10548408.2015.1038418>
- Angerschmid, A., Zhou, J., Theuermann, K., Chen, F., & Holzinger, A. (2022). Fairness and Explanation in AI-Informed Decision Making. *Machine Learning and Knowledge Extraction*, 4(2), 556–579. <https://doi.org/10.3390/make4020026>
- Begu, T. M., Soica, S., & Nedelcu, A. (2025). The Impact of AI on Digital Quality and Technical Sustainability of Travel Websites. *Sustainability*, 17(21), 9879. <https://doi.org/10.3390/su17219879>
- Bhardwaj, S., Jain, V., Mahapatra, D., & Sindhwani, R. (2025). Exploring the Dark Side of AI and Its Influence on Consumer Emotion. *Journal of Consumer Behaviour*, 24(2), 529–544. <https://doi.org/10.1002/cb.2431>

- Bietti, E. (2021). From Ethics Washing to Ethics Bashing: A Moral Philosophy View on Tech Ethics. *Journal of Social Computing*, 2(3), 266–283. <https://doi.org/10.23919/JSC.2021.0031>
- Boksberger, P. E., & Melsen, L. (2011). Perceived value: A critical examination of definitions, concepts and measures for the service industry. *Journal of Services Marketing*, 25(3), 229–240. <https://doi.org/10.1108/08876041111129209>
- Brace, I. (2005). *Questionnaire design: How to plan, structure and write survey material for effective market research*. Kogan Page.
- Bramer, M., & Stahl, F. (Eds.). (2025). *Artificial Intelligence XLII: 45th SGAI International Conference on Artificial Intelligence, AI 2025, Cambridge, UK, December 16-18, 2025, Proceedings, Part II* (Vol. 16302). Springer Nature Switzerland. <https://doi.org/10.1007/978-3-032-11442-6>
- Chen, C., & Wei, Z. (2024). Role of Artificial Intelligence in travel decision making and tourism product selling. *Asia Pacific Journal of Tourism Research*, 29(3), 239–253. <https://doi.org/10.1080/10941665.2024.2317390>
- Choung, H., David, P., & Ross, A. (2023). Trust and ethics in AI. *AI & SOCIETY*, 38(2), 733–745. <https://doi.org/10.1007/s00146-022-01473-4>
- Cox, A. (2024). Algorithmic Literacy, AI Literacy and Responsible Generative AI Literacy. *Journal of Web Librarianship*, 18(3), 93–110. <https://doi.org/10.1080/19322909.2024.2395341>
- Cristian, M. G., Tileagă, C., Lucian Blaga University of Sibiu, Romania, & Lucian Blaga University of Sibiu, Romania. (2024). CHALLENGES AND PERSPECTIVES OF AI IN SUSTAINABLE TOURISM. *Management of Sustainable Development*, 16(2), 14–26. <https://doi.org/10.54989/msd-2024-0012>
- Dang, L., Steffen, A., Weibel, C., & Von Arx, W. (2023). Behavioural pricing effects in tourism: A review of the empirical evidence and its managerial implications. *European Journal of Tourism Research*, 36, 3603. <https://doi.org/10.54055/ejtr.v36i.2850>
- Del Valle, J. I., & Lara, F. (2024). AI-powered recommender systems and the preservation of personal autonomy. *AI & SOCIETY*, 39(5), 2479–2491. <https://doi.org/10.1007/s00146-023-01720-2>

- Dogruel, L., Masur, P., & Joeckel, S. (2022). Development and Validation of an Algorithm Literacy Scale for Internet Users. *Communication Methods and Measures*, 16(2), 115–133.
<https://doi.org/10.1080/19312458.2021.1968361>
- Fanti, L., Guarascio, D., & Moggi, M. (2022). From Heron of Alexandria to Amazon's Alexa: A stylized history of AI and its impact on business models, organization and work. *Journal of Industrial and Business Economics*, 49(3), 409–440. <https://doi.org/10.1007/s40812-022-00222-4>
- Faul, F., Erdfelder, E., Lang, A.-G., & Buchner, A. (2007). G*Power 3: A flexible statistical power analysis program for the social, behavioral, and biomedical sciences. *Behavior Research Methods*, 39(2), 175–191. <https://doi.org/10.3758/BF03193146>
- Ferrer-Rosell, B., Massimo, D., & Berezina, K. (Eds.). (2023). *Information and Communication Technologies in Tourism 2023: Proceedings of the ENTER 2023 eTourism Conference, January 18-20, 2023*. Springer Nature Switzerland. <https://doi.org/10.1007/978-3-031-25752-0>
- Filieri, R., Alguezaui, S., & McLeay, F. (2015). Why do travelers trust TripAdvisor? Antecedents of trust towards consumer-generated media and its influence on recommendation adoption and word of mouth. *Tourism Management*, 51, 174–185.
<https://doi.org/10.1016/j.tourman.2015.05.007>
- Gagrčin, E., Naab, T. K., & Grub, M. F. (2024). Algorithmic media use and algorithm literacy: An integrative literature review. *New Media & Society*, 28(1), 423–447.
<https://doi.org/10.1177/14614448241291137>
- Gil De Zúñiga, H., Goyanes, M., & Durotoye, T. (2024). A Scholarly Definition of Artificial Intelligence (AI): Advancing AI as a Conceptual Framework in Communication Research. *Political Communication*, 41(2), 317–334. <https://doi.org/10.1080/10584609.2023.2290497>
- González-Sanz, E. L., Cantador, I., & Bellogín, A. (2025). *LLM-based Generation of Personalized, Context-aware City Tourist Itineraries: A User Study with GPT Trip Planner*.
- Gretzel, U., Sigala, M., Xiang, Z., & Koo, C. (2015). Smart tourism: Foundations and developments. *Electronic Markets*, 25(3), 179–188. <https://doi.org/10.1007/s12525-015-0196-8>

- Gu, S. (2024). A Survey of Large Language Models in Tourism (Tourism LLMs). *Qeios*.
<https://doi.org/10.32388/8R27CJ>
- Haenlein, M., & Kaplan, A. (2019). A Brief History of Artificial Intelligence: On the Past, Present, and Future of Artificial Intelligence. *California Management Review*, 61(4), 5–14.
<https://doi.org/10.1177/0008125619864925>
- Hassan, N., Abdelraouf, M., & El-Shihy, D. (2025). The moderating role of personalized recommendations in the trust–satisfaction–loyalty relationship: An empirical study of AI-driven e-commerce. *Future Business Journal*, 11(1), 66. <https://doi.org/10.1186/s43093-025-00476-z>
- Hayes, A. F. (2022). *Introduction to mediation, moderation, and conditional process analysis: A regression-based approach* (Third edition). The Guilford Press.
- Hlee, S., Yan, Z., & Li, P. (2025). AI-powered travel recommendations and decision-making: The role of spatio-temporal efficiency, destination type, and travel party composition. *Electronic Markets*, 35(1), 75. <https://doi.org/10.1007/s12525-025-00825-4>
- Hultman, M., Kazeminia, A., & Ghasemi, V. (2015). Intention to visit and willingness to pay premium for ecotourism: The impact of attitude, materialism, and motivation. *Journal of Business Research*, 68(9), 1854–1861. <https://doi.org/10.1016/j.jbusres.2015.01.013>
- Igartua, J.-J., & Hayes, A. F. (2021). Mediation, Moderation, and Conditional Process Analysis: Concepts, Computations, and Some Common Confusions. *The Spanish Journal of Psychology*, 24, e49. <https://doi.org/10.1017/SJP.2021.46>
- Ihrke, M. (2011). Automatic generation of randomized trial sequences for priming experiments. *Frontiers in Psychology*, 2. <https://doi.org/10.3389/fpsyg.2011.00225>
- Janiesch, C., Zschech, P., & Heinrich, K. (2021). Machine learning and deep learning. *Electronic Markets*, 31(3), 685–695. <https://doi.org/10.1007/s12525-021-00475-2>
- Jannach, D., Naveed, S., & Jugovac, M. (2017). User Control in Recommender Systems: Overview and Interaction Challenges. In D. Bridge & H. Stuckenschmidt (Eds.), *E-Commerce and Web Technologies* (Vol. 278, pp. 21–33). Springer International Publishing.
https://doi.org/10.1007/978-3-319-53676-7_2

- Jasso, G. (2006). Factorial Survey Methods for Studying Beliefs and Judgments. *Sociological Methods & Research*, 34(3), 334–423. <https://doi.org/10.1177/0049124105283121>
- Jha, G. K., Gaur, M., Ranjan, P., & Thakur, H. K. (2023). A survey on trustworthy model of recommender system. *International Journal of System Assurance Engineering and Management*, 14(S3), 789–806. <https://doi.org/10.1007/s13198-021-01085-z>
- Jiang, W., Li, D., & Liu, C. (2025). Understanding dimensions of trust in AI through quantitative cognition: Implications for human-AI collaboration. *PLOS One*, 20(7), e0326558. <https://doi.org/10.1371/journal.pone.0326558>
- Kazim, E., & Koshiyama, A. S. (2021). A high-level overview of AI ethics. *Patterns*, 2(9), 100314. <https://doi.org/10.1016/j.patter.2021.100314>
- Kim, K., & Ha, H.-Y. (2024). The dynamics of perceived justice and its outcomes in the online tourism sector: Inter-relationships and temporal and carryover effects. *Current Issues in Tourism*, 27(24), 4740–4756. <https://doi.org/10.1080/13683500.2023.2280145>
- Kim, S. Y., & Kim, J. (2025). The impact of AI recommendation quality on service satisfaction: The moderating roles of standardization and customization. *Journal of Services Marketing*, 39(4), 365–386. <https://doi.org/10.1108/JSM-05-2024-0214>
- Kontogianni, A., Alepis, E., Virvou, M., & Patsakis, C. (2024). *Smart Tourism—The Impact of Artificial Intelligence and Blockchain* (Vol. 249). Springer Nature Switzerland. <https://doi.org/10.1007/978-3-031-50883-7>
- Koo, I., Zaman, U., Ha, H., & Nawaz, S. (2025). Assessing the interplay of trust dynamics, personalization, ethical AI practices, and tourist behavior in the adoption of AI-driven smart tourism technologies. *Journal of Open Innovation: Technology, Market, and Complexity*, 11(1), 100455. <https://doi.org/10.1016/j.joitmc.2024.100455>
- Koo, M., & Yang, S.-W. (2025). Likert-Type Scale. *Encyclopedia*, 5(1), 18. <https://doi.org/10.3390/encyclopedia5010018>
- Kopalle, P. K., Gangwar, M., Kaplan, A., Ramachandran, D., Reinartz, W., & Rindfleisch, A. (2022). Examining artificial intelligence (AI) technologies in marketing via a global lens: Current

- trends and future research opportunities. *International Journal of Research in Marketing*, 39(2), 522–540. <https://doi.org/10.1016/j.ijresmar.2021.11.002>
- Li, C.-H. (2016). Confirmatory factor analysis with ordinal data: Comparing robust maximum likelihood and diagonally weighted least squares. *Behavior Research Methods*, 48(3), 936–949. <https://doi.org/10.3758/s13428-015-0619-7>
- Li, M., Yin, D., Qiu, H., & Bai, B. (2021). A systematic review of AI technology-based service encounters: Implications for hospitality and tourism operations. *International Journal of Hospitality Management*, 95, 102930. <https://doi.org/10.1016/j.ijhm.2021.102930>
- Li, Y., Wu, B., Huang, Y., & Luan, S. (2024). Developing trustworthy artificial intelligence: Insights from research on interpersonal, human-automation, and human-AI trust. *Frontiers in Psychology*, 15, 1382693. <https://doi.org/10.3389/fpsyg.2024.1382693>
- Lim, S. (Edward), & Kim, M. (2025). AI-powered personalized recommendations and pricing: Moderating effects of ethical AI and consumer empowerment. *International Journal of Hospitality Management*, 130, 104259. <https://doi.org/10.1016/j.ijhm.2025.104259>
- Lin, W., Zhou, X., Sun, L., Qi, L., Tsai, S.-B., Yang, Y., Liu, H., Kou, H., & Kong, L. (2025). A Locality-Sensitive-Hashing-Based Collaborative Recommendation Method for Responsible AI-Driven Recommender Systems. *IEEE Transactions on Artificial Intelligence*, 6(2), 393–404. <https://doi.org/10.1109/TAI.2024.3381571>
- Long, X., Sun, J., Dai, H., Zhang, D., Zhang, J., Chen, Y., Hu, H., & Zhao, B. (2025). The Choice Overload Effect in Online Recommender Systems. *Manufacturing & Service Operations Management*, 27(1), 249–268. <https://doi.org/10.1287/msom.2022.0659>
- Lu, W. (2024). Inevitable challenges of autonomy: Ethical concerns in personalized algorithmic decision-making. *Humanities and Social Sciences Communications*, 11(1), 1321. <https://doi.org/10.1057/s41599-024-03864-y>
- Luo, Y., Kumar, N., & Yazdanmehr, A. (2024). *AI Nudging and Decision Quality: Evidence from Randomized Experiments in Online Recommendation Setting*. SSRN. <https://doi.org/10.2139/ssrn.4967812>

- Lv, L., Chen, S., Liu, G. G., & Benckendorff, P. (2024). Enhancing customers' life satisfaction through AI-powered personalized luxury recommendations in luxury tourism marketing. *International Journal of Hospitality Management*, 123, 103914.
<https://doi.org/10.1016/j.ijhm.2024.103914>
- MacKinnon, D. P., Krull, J. L., & Lockwood, C. M. (2000). *Equivalence of the Mediation, Confounding and Suppression Effect*. 1 (4), 173–181.
<https://doi.org/https://doi.org/10.1023/A:1026595011371>
- Malhotra, N. K., Nunan, D., & Birks, D. F. (2017). *Marketing research: An applied approach* (Fifth edition). Pearson.
- McDermid, J. A., Jia, Y., Porter, Z., & Habli, I. (2021). Artificial intelligence explainability: The technical and ethical dimensions. *Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences*, 379(2207), 20200363.
<https://doi.org/10.1098/rsta.2020.0363>
- Meade, A. W., & Craig, S. B. (2012). Identifying careless responses in survey data. *Psychological Methods*, 17(3), 437–455. <https://doi.org/10.1037/a0028085>
- Mehrabian, A., & Russell, J. A. (1974). *An approach to environmental psychology*. MIT Press.
- Milano, S., Taddeo, M., & Floridi, L. (2020). Recommender systems and their ethical challenges. *AI & SOCIETY*, 35(4), 957–967. <https://doi.org/10.1007/s00146-020-00950-y>
- Milwood, P. A., Hartman-Caverly, S., & Roehl, W. S. (2023). A Scoping Study of Ethics in Artificial Intelligence Research in Tourism and Hospitality. In B. Ferrer-Rosell, D. Massimo, & K. Berezina (Eds.), *Information and Communication Technologies in Tourism 2023* (pp. 243–254). Springer Nature Switzerland. https://doi.org/10.1007/978-3-031-25752-0_26
- Nguyen, M. (2025). *Experimental Design*. Springer Nature Switzerland. <https://doi.org/10.1007/978-3-032-01839-7>
- Nizamani, M. M., Zhang, H.-L., & Lai, Z. (2025). Human-centered AI: Advancing ethical, transparent, and context-aware systems for sustainable development. *Technology in Society*, 84, 103121. <https://doi.org/10.1016/j.techsoc.2025.103121>
- Nunnally, J. C. (1978). *Psychometric Theory*. (2nd ed.). McGraw-Hill.

- Oeldorf-Hirsch, A., & Neubaum, G. (2023). *What do we know about algorithmic literacy? The status quo and a research agenda for a growing field*. 27(2), 681–701.
<https://doi.org/10.1177/14614448231182662>
- Oliver, R. L. (1980). A Cognitive Model of the Antecedents and Consequences of Satisfaction Decisions. *Journal of Marketing Research*, 17(4), 460. <https://doi.org/10.2307/3150499>
- Oliveri, A. M., Polizzi, G., & Parroco, A. M. (2019). Measuring Tourist Satisfaction Through a Dual Approach: The 4Q Methodology. *Social Indicators Research*, 146(1–2), 361–382.
<https://doi.org/10.1007/s11205-018-2013-1>
- Onghena, P. (2017). Randomization and the Randomization Test: Two Sides of the Same Coin. In V. W. Berger (Ed.), *Randomization, Masking, and Allocation Concealment* (1st ed., pp. 185–208). Chapman and Hall/CRC. <https://doi.org/10.1201/9781315305110-13>
- Pacherie, E. (2007). *The Sense of Control and the Sense of Agency: Psyche*. 13 (1), 1–30.
- Parasuraman, R., & Manzey, D. H. (2010). Complacency and Bias in Human Use of Automation: An Attentional Integration. *Human Factors: The Journal of the Human Factors and Ergonomics Society*, 52(3), 381–410. <https://doi.org/10.1177/0018720810376055>
- Park, Y. E., & Son, H. (2025). AI recommendation vs. crowdsourced recommendation vs. travel expert recommendation: The moderating role of consumption goal on travel destination decision. *PLOS ONE*, 20(3), e0318719. <https://doi.org/10.1371/journal.pone.0318719>
- Pavlou, P. A. (2003). Consumer Acceptance of Electronic Commerce: Integrating Trust and Risk with the Technology Acceptance Model. *International Journal of Electronic Commerce*, 7(3), 101–134. <https://doi.org/10.1080/10864415.2003.11044275>
- Podsakoff, P. M., MacKenzie, S. B., Lee, J.-Y., & Podsakoff, N. P. (2003). Common method biases in behavioral research: A critical review of the literature and recommended remedies. *Journal of Applied Psychology*, 88(5), 879–903. <https://doi.org/10.1037/0021-9010.88.5.879>
- Raees, M., Meijerink, I., Lykourantzou, I., Khan, V.-J., & Papangelis, K. (2024). From explainable to interactive AI: A literature review on current trends in human-AI interaction. *International Journal of Human-Computer Studies*, 189, 103301.
<https://doi.org/10.1016/j.ijhcs.2024.103301>

- Ramani, R. S., & Aguinis, H. (2023). Using field and quasi experiments and text-based analysis to advance international business theory. *Journal of World Business*, 58(5), 101463.
<https://doi.org/10.1016/j.jwb.2023.101463>
- Ranjbaran, A., Shabankareh, M., Seyyedamiri, N., & Jandaghi, G. (2024). The concept of the new tourism destinations: An exploration of satisfaction and dissatisfaction factors through visitors' reviews. *Journal of Organisational Studies and Innovation*, 11(2), 1–18.
<https://doi.org/10.51659/josi.23.190>
- Riquelme, I. P., Román, S., & Cuestas, P. J. (2021). Does it matter who gets a better price? Antecedents and consequences of online price unfairness for advantaged and disadvantaged consumers. *Tourism Management Perspectives*, 40, 100902.
<https://doi.org/10.1016/j.tmp.2021.100902>
- Rubio Juan, M., & Revilla, M. (2021). Comparing respondents who passed versus failed an Instructional Manipulation Check: A case study about support for climate change policies. *International Journal of Market Research*, 63(4), 408–415.
<https://doi.org/10.1177/14707853211023039>
- Ryan, R. M., & Deci, E. L. (2020). Intrinsic and extrinsic motivation from a self-determination theory perspective: Definitions, theory, practices, and future directions. *Contemporary Educational Psychology*, 61, 101860. <https://doi.org/10.1016/j.cedpsych.2020.101860>
- Samala, N., Katkam, B. S., Bellamkonda, R. S., & Rodriguez, R. V. (2022). Impact of AI and robotics in the tourism sector: A critical insight. *Journal of Tourism Futures*, 8(1), 73–87.
<https://doi.org/10.1108/JTF-07-2019-0065>
- Santosh, K., & Wall, C. (2022). *AI, Ethical Issues and Explainability—Applied Biometrics*. Springer Nature Singapore. <https://doi.org/10.1007/978-981-19-3935-8>
- Sarkar, J. L., Majumder, A., Panigrahi, C. R., Roy, S., & Pati, B. (2023). Tourism recommendation system: A survey and future research directions. *Multimedia Tools and Applications*, 82(6), 8983–9027. <https://doi.org/10.1007/s11042-022-12167-w>
- Siegel, G. (2025). The Perceptual Sense of Agency. *Review of Philosophy and Psychology*, 16(4), 1473–1499. <https://doi.org/10.1007/s13164-025-00788-7>

- Singu, H. B., Chakraborty, D., Troise, C., Camilleri, M. A., & Bresciani, S. (2025). Responsible AI for trustworthy tourism: A framework for mitigating ambiguity and anxiety with generative AI. *Technological Forecasting and Social Change*, 223, 124407.
<https://doi.org/10.1016/j.techfore.2025.124407>
- Solano-Barliza, A., Arregocés-Julio, I., Aarón-Gonzalvez, M., Zamora-Musa, R., De-La-Hoz-Franco, E., Escorcía-Gutierrez, J., & Acosta-Coll, M. (2024). Recommender systems applied to the tourism industry: A literature review. *Cogent Business & Management*, 11(1), 2367088.
<https://doi.org/10.1080/23311975.2024.2367088>
- Spence, M. (1973). Job Market Signaling. *The Quarterly Journal of Economics*, 87(3), 355.
<https://doi.org/10.2307/1882010>
- Sullivan, G. M., & Feinn, R. (2012). Using Effect Size—Or Why the P Value Is Not Enough. *Journal of Graduate Medical Education*, 4(3), 279–282. <https://doi.org/10.4300/JGME-D-12-00156.1>
- Sun, Z. Q., & Yoon, S. J. (2022). What Makes People Pay Premium Price for Eco-Friendly Products? The Effects of Ethical Consumption Consciousness, CSR, and Product Quality. *Sustainability*, 14(23), 15513. <https://doi.org/10.3390/su142315513>
- Twabu, K. (2025). Enhancing the cognitive load theory and multimedia learning framework with AI insight. *Discover Education*, 4(1), 160. <https://doi.org/10.1007/s44217-025-00592-6>
- Vasiliki, M., Serdaris, P., Antoniadis, I., & Spinthiropoulos, K. (2025). Ethical Consumer Attitudes and Trust in Artificial Intelligence in the Digital Marketplace: An Empirical Analysis of Behavioral and Value-Driven Determinants. *Digital*, 6(1), 1.
<https://doi.org/10.3390/digital6010001>
- Walek, B., & Sládek, O. (2025). Comparison of Selected Algorithms in Movie Recommender System. *Applied Sciences*, 15(17), 9518. <https://doi.org/10.3390/app15179518>
- Wan, J., & Bharti, M. (2025). From Price to Quantity: Redefining How Consumers Pay the Ethical Premium. *Journal of Business Ethics*. <https://doi.org/10.1007/s10551-025-06178-4>
- Wang, K., Lu, L., Fang, J., Xing, Y., Tong, Z., & Wang, L. (2023). The downside of artificial intelligence (AI) in green choices: How AI recommender systems decrease green

- consumption. *Managerial and Decision Economics*, 44(6), 3346–3353.
<https://doi.org/10.1002/mde.3882>
- Wang, R., Wu, C., Wang, X., Xu, F., & Yuan, Q. (2024). e-Tourism Information Literacy and Its Role in Driving Tourist Satisfaction With Online Travel Information: A Qualitative Comparative Analysis. *Journal of Travel Research*, 63(4), 904–922.
<https://doi.org/10.1177/00472875231177229>
- Woodgate, J., & Ajmeri, N. (2024). Macro Ethics Principles for Responsible AI Systems: Taxonomy and Directions. *ACM Computing Surveys*, 56(11), 1–37. <https://doi.org/10.1145/3672394>
- Wu, F., & Xie, G. (2024). Research on the influence of perceived quality on users' continuance usage intention of online live streaming class platforms: The mediating role of flow experience and the moderating impact of perceived usefulness. *Frontiers in Psychology*, 15, 1455597.
<https://doi.org/10.3389/fpsyg.2024.1455597>
- Xu, Y., Chen, Z., & Dong, M. (2025). Shaping the fairness journey: The roles of AI literacy, explanation, and interpersonal interaction in AI interviews. *International Journal of Human-Computer Studies*, 205, 103629. <https://doi.org/10.1016/j.ijhcs.2025.103629>
- Yang, H., Li, D., & Hu, P. (2024). Decoding algorithm fatigue: The role of algorithmic literacy, information cocoons, and algorithmic opacity. *Technology in Society*, 79, 102749.
<https://doi.org/10.1016/j.techsoc.2024.102749>
- Yu, X., Yang, Y., & Li, S. (2024). Users' continuance intention towards an AI painting application: An extended expectation confirmation model. *PLOS ONE*, 19(5), e0301821.
<https://doi.org/10.1371/journal.pone.0301821>
- Zhang, B., & Sundar, S. S. (2019). Proactive vs. reactive personalization: Can customization of privacy enhance user experience? *International Journal of Human-Computer Studies*, 128, 86–99. <https://doi.org/10.1016/j.ijhcs.2019.03.002>
- Zhao, Z., Fan, W., Li, J., Liu, Y., Mei, X., Wang, Y., Wen, Z., Wang, F., Zhao, X., Tang, J., & Li, Q. (2024). Recommender Systems in the Era of Large Language Models (LLMs). *IEEE Transactions on Knowledge and Data Engineering*, 36(11), 6889–6907.
<https://doi.org/10.1109/TKDE.2024.3392335>

EXECUTIVE SUMMARY

This thesis explores whether giving travellers more control over the ethical settings of AI recommendation systems can improve their satisfaction. Using a between-subjects experiment with 122 participants, respondents were presented with either a standard Booking.com-style interface or a version that included visible ethical preference controls related to eco-friendliness, data privacy, and social impact.

The findings showed that the ethical sliders themselves did not directly increase satisfaction. Instead, their effect came from the greater sense of control they provided to the users. Perceived agency emerged as a significant mediator, suggesting that participants responded more positively when they felt actively involved in shaping the recommendation process. In contrast, platform trust and perceived ethical alignment did not significantly mediate the relationship, possibly because these perceptions develop over time through repeated interactions rather than from a single exposure. Additionally, exploratory factor analysis revealed limited discriminant validity between these three constructs in a single-exposure context.

The study also found that users with higher levels of algorithmic literacy were generally more satisfied with AI-powered recommendation systems, although they also appeared more critical of superficial ethical signs. Overall, the results suggest that ethical customisation features are most valuable when they enhance users' sense of autonomy rather than simply serving as ethical branding.

For platform designers, the central implication is that making ethical controls visible and clearly labelled, even before users interact with them, is a low-cost intervention with meaningful psychological impact.

KEYWORDS: Artificial Intelligence (AI), Ethical AI, Recommender Systems, Traveller Satisfaction, Perceived Agency, Platform Trust, Algorithmic Literacy, Human-Centred AI, Digital Tourism

WORD COUNT: 19 980



Ecole de Gestion de l'Université de Liège