

Article

Not peer-reviewed version

Decoding the Trust–Tech Nexus in AI-Driven Academic Marketing: A Multi-Paradigm Machine Learning Analysis of Stakeholder Decision-Making Confidence

[Pradnya Dalavi](#)*, [Ganesh Waghmare](#), [Ravindra Khedkar](#)

Posted Date: 10 February 2026

doi: 10.20944/preprints202602.0775.v1

Keywords: artificial intelligence; academic marketing; stakeholder trust; machine learning; decision-making confidence; trust–tech nexus; higher education



Preprints.org is a free multidisciplinary platform providing preprint service that is dedicated to making early versions of research outputs permanently available and citable. Preprints posted at Preprints.org appear in Web of Science, Crossref, Google Scholar, Scilit, Europe PMC.

Copyright: This open access article is published under a [Creative Commons CC BY 4.0 license](#), which permit the free download, distribution, and reuse, provided that the author and preprint are cited in any reuse.

Disclaimer/Publisher's Note: The statements, opinions, and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions, or products referred to in the content.

Article

Decoding the Trust–Tech Nexus in AI-Driven Academic Marketing: A Multi-Paradigm Machine Learning Analysis of Stakeholder Decision-Making Confidence

Pradnya Dalavi *, Ganesh Waghmare and Ravindra Khedkar

MIT-ADT University, Loni Kalbhor, Pune, India

* Correspondence: researchpradnya@gmail.com

Abstract

The growing presence of AI in academic marketing is reshaping the way HEIs (Higher Education Institutions) are engaging their stakeholders. The benefits that come from utilizing AI-driven elements to improve operational efficiency only come to fruition when stakeholders trust those integrated elements. The relationship between stakeholder perceptions and their confidence in decision making on AI-enabled marketing activities is examined by this study through the application of a multi-paradigm machine-learning framework. This study analyzed the responses of 200 stakeholders using a combination of survey data and provided greater strength to the findings by combining multiple responses. Five different paradigms were used: Linear Regression, Ridge Regression, Support Vector Regression, Random Forest, and Gradient Boosting to evaluate stakeholder data and predict their behavior. The results suggest that Ridge Regression provided the most stable baseline for prediction; however, the ensemble models were able to capture many critical non-linear dynamics that were overlooked by linear model approaches. Trust in AI tools and personalised advertising is a dominant factor in stakeholder confidence, whereas the institutional perception of stakeholders acts as a structural moderator. These findings provide empirical validation of the Trust-Tech Nexus and demonstrate that for stakeholders to fully realise the benefits of AI-driven personalised advertising, institutional credibility is essential.

Keywords: artificial intelligence; academic marketing; stakeholder trust; machine learning; decision-making confidence; trust–tech nexus; higher education

1. Introduction

Rapid advancements in Artificial Intelligence (AI) are changing the way digital marketing is done. Data-driven personalization, predictive analytics, and Automated Decision Support (ADS) are being used at scale for digital marketing. AI tools, such as personalized ads, recommendation engines and chatbot technology, are being adopted across higher education ecosystems to engage prospective students and other stakeholders. The use of these technologies will make the process of communication more relevant, efficient, and responsive; how effective they are rests with how stakeholders perceive and trust the use of AI in engaging them [1,2]. Higher Education Institutions (HEIs) operate in a context of high choice and high-risk decision-making; stakeholder decisions such as what program to enroll in or what institution to attend will have long-term consequences both personally, professionally, and financially. Past studies based on the Technology Acceptance Model (TAM) recognize that perceived usefulness and ease of use of a technology are primary factors influencing whether a user will adopt a technology [3,4] While TAM provides a good basis for understanding the factors influencing technology adoption, recent studies suggest that it does not fully capture the complexity of AI-enabled environments, as the algorithmic autonomy, opacity, and

adaptive personalization of AI introduce three new areas of perceived risk, control and accountability [5]. Previous studies exploring AI and digital marketing consistently indicate that trust plays an important role in a user’s decision to accept a new technology, and therefore trust will shape user attitudes toward AI-powered (algorithmic) systems, beyond the functional performance of the technology [6,7]. In the context of AI-powered marketing, there is often a privacy-personalization paradox, where stakeholders often value personalized content but also express concerns about how their data is used, monitored (surveillance) and the implications of automated decision-making systems [8]. The challenges associated with trust in the realm of higher education also includes institutional credibility as ethical stewards, and reputation integrity impacts stakeholder confidence and judgments of legitimacy [9]. Importantly, trust in AI does not exist in isolation. Relationship marketing theory suggests that the level of attributed institutional trust will influence how stakeholders interpret and evaluate institutional actions[10]. In the context of higher education with institutional marketing, this means that although technically advanced AI systems may improve a user’s confidence, they may not achieve this goal if the institution does not demonstrate a level of transparency and ethical responsibility. Conversely, if the institution has a strong perception from its stakeholders, this perception will have an amplifying effect on the positive aspects of AI-powered personalization and recommendations. Although previous studies have acknowledged the importance of institutional trust in a digital environment, the relationship between institutional trust and AI is still being studied using a small number of empirical data-driven approaches [5]. From a methodological perspective, most of the current literature has used linear regression-based techniques which assume that the effect of AI is additive and homogeneous across all users. However, as the existing literature has noted, the way humans respond to AI is often dependent on a variety of factors and is often non-linear, and threshold dependent [11–13]. Machine learning (ML) and specifically ensemble learning approaches may allow you to model such complex behavioral reactions to AI effectively, however, ML techniques are often criticized for the lack of interpretability of their findings in the context of social sciences research [14]. Addressing this long-standing tension between predictability and theoretical grounding requires the integration of predictive modeling with theoretical and explainable analytical frameworks.

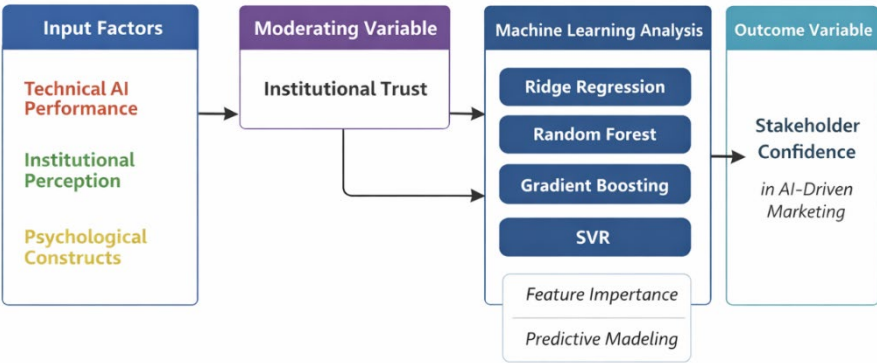


Figure 1. Conceptual framework illustrating the Trust–Tech Nexus in AI-driven academic marketing.

This research uses a multi-method ‘Trust-Tech Nexus’ conceptual framework informed by Copeland et al., 2020 and based on Learning Management Systems (LMS) and Technology Acceptance Models (TAM). The conceptual framework identifies the key interactions between AI technical attributes and other factors such as; trust in AI tools; perception of the University; and emotional reactions and skepticism about artificial intelligence; influencing the level of confidence decision-makers have when using AI-driven methods of marketing. In addition, the conceptual framework provides insight into the structure of the study; including variable inclusion and analytical techniques employed. To empirically test this conceptual framework, we take a multi-faceted machine learning approach. The study incorporates linear regularized modelling, kernel

modelling, ensemble modelling and predictive stable modelling to create predictive models, which capture linear and non-linear relationships and interactions between the variables. We utilize a machine learning approach alongside established behavioral theory to enhance the understanding of AI confidence and therefore the development of policies and practices by HEIs that use artificial intelligence to improve the level of confidence decision-makers have when using AI for marketing purposes. There are three main contributions; We empirically validate trust-based interpretations of technology acceptance, develop an understanding of the value of combining explainable machine learning and established behavioral theories, and provide actionable insights for Higher Education Institutions (HEIs) to align their AI-driven personalised marketing with their trust-building strategies in relation to their institutions.

2. Materials and Methods

In this study, we present a multi-stage predictive analytics framework that incorporates traditional psychometric survey methods combined with modern machine learning (ML) techniques. Unlike earlier methods where only statistical inferential analyses are conducted, we will adopt an expanded view of predictive validity, providing more emphasis on theoretical constructs being able to generalize to previously unseen data[15,16]. Predictive validity paradigms are most appropriate for investigating complex behavioral constructs such as confidence among stakeholders to make decisions [2] given that the relationships are anticipated to be non-linear, interplay-driven and dependent on the context.

There are four distinct phases or stages in this methodological pipeline:

- (1) collecting primary data,
- (2) augmenting primary data while preserving covariance,
- (3) creating models across multiple paradigms,
- (4) analyzing robustness and interpretability.

2.1. Research Design and Data Collection

In the study, primary data were gathered using a cross-sectional survey design from individuals or groups who had previous experiences using AI-enabled tools to Market Academic Programs through universities, such as personalized advertising, recommendation systems, and chatbots. This sort of design is similar to how prior researchers have studied user perceptions of and trust in AI-based systems [17–19]. To gather the data, the authors developed a survey based partially on existing scales that had verified reliability and validity through the Technology Acceptance Model [20,21] and Trust-Commitment Theory [10]. In the end, they looked at ten latent constructs, which included not only technical features of AI tools (like predictive recommendations, chatbot responsiveness, and personalized service), but also psychological and organizational factors that contact with the tool could engender (such as trusting in AI-based systems, perceptions of colleges/universities, emotional involvement with AI, and skepticism regarding AI). All survey items were evaluated on a 5-point Likert scale from “Strongly Disagree” (1) to “Strongly Agree” (5). This sort of option for response is common in trust and health technology research and provides researchers with statistically significant data while offering reasonable interpretive power for survey participants[17,22,23]. After removing all incomplete surveys and outliers, a total of 200 surveys were available to analyze. Although this total is appropriate for standard regression methods, research on methodology indicates that ensemble learning models trained on relatively small data sets may generate models with higher errors [24–26] Therefore, additional steps were taken to ensure stable and generalizable models.

2.2. Data Augmentation via Covariance-Preserving Sampling

The effective sample size was increased to 500 through the use of Multivariate Normal Distribution (MVN) sampling techniques to reduce the impact of sampling variability and instability

of estimators. In particular, the MVN sampling technique generated an additional 300 synthetic observations from the original data using the empirical mean vector (μ) and covariance matrix (Σ). In contrast to simple oversampling and noise augmentation methods, the MVN sampling technique preserves the covariance structure between the variables, thus preserving the meaningful theoretical relationships between variables (e.g., the relationship between institutional perception and decision-making confidence), which was maintained in the augmented data. Covariances are particularly important in psychometric research, where the constructs examined are interdependent rather than statistically independent [27,28]. The augmented dataset will be used solely for predictive modeling, cross-validation, and robustness analysis. That is, no causal inference will be made from the augmented data. This is congruent with current best practices in applying machine learning techniques within the social and behavioral sciences [29,30].

2.3. Machine Learning Models and Benchmarking Strategy

The authors adopted a multi-paradigm modeling approach to assess different types of learning models (linear, regularized, kernel-based, and ensemble). With their benchmarking models, the authors would explicitly determine if stakeholder confidence in making decisions is more dependent on linear associations or on nonlinear and higher-order associations, which is a common characteristic of trust-related behavior. The authors chose the ordinary least squares (OLS) regression model as the benchmark model for their study. The OLS regression model served as a standard for comparison with other statistical modeling methods by providing an interpretable set of coefficients. Additionally, to prevent the negative effects of multicollinearity on the coefficients from OLS regression, the authors utilized a regularization method called ridge regression [25,31]. The authors used support vector regression (SVR) as the second type of modeling used in this study. SVR utilized a radial basis function (RBF) kernel to examine non-linear relationships among the bounded Likert-type scale data, while offering the additional advantage of noise robustness [32,33]. For capturing complex interaction effects between variables in this study, the authors utilized ensemble learning approaches. For example, in their research, the authors applied Random Forest [34,35] as a bagged tree ensemble approach to analyze non-linear, threshold effects and higher-order interactions. Random Forest is also helpful to decrease variance through the average of many trees. The authors used Gradient Boosting Machine (GBM), which is a boosting-style ensemble method that uses residual errors from the previous prediction to improve future predictions, as a means to further enhance prediction accuracy. One of the primary advantages of GBM is its ability to model subtle, high-order, non-linear relationships [36,37]. Based on their findings, the authors concluded that using a multi-paradigm strategy is ideal in that it allows for simultaneous assessment of the three aspects of predictive performance, behaviorally interpretable models, and methodologically sound modelling methods, rather than optimally assessing these aspects separately.

2.4. Model Validation and Performance Evaluation

Cross-validation with 10 folds allowed for the calculation of performance estimates for a model to ensure bias-free performance estimates. The data set was separated into 10 parts and each part was designated as the testing set during the training of the model; 9 parts were used to train the model and 1 part was used for testing during each training run until all parts had been utilized for both tests and training. Performance of the predictive models was evaluated using Root Mean Square Error (RMSE) as the primary performance measure, while R^2 was a secondary performance measure. When a validation method such as 10-fold cross-validation is used, it reduces the likelihood that a single split of the data will significantly affect overall model performance and allows for the ability to determine if model performance is truly reflective of generalizable modelling ability instead of being based on random data splits [38,39].

2.5. Robustness Testing and Model Interpretability

Bootstrapping (1,000 iterations) enables one to generate a 95% confidence interval around the RMSE and R² values of their model by bootstrapping these metrics. This will provide insight on how sensitive the models are with respect to sampling error, thus increasing the confidence of the results presented. Ensemble learning models may also have limited interpretability, as it is difficult to determine the individual predictors and the contribution of each; therefore, a Permutation Feature Importance methodology has been applied to provide an agnostic explanation of ensemble models. This methodology quantifies how much each predictor contributes to the predictive power/performance of each model by applying random permutations of each predictor variable within each iteration. Because many behavioral datasets are anonymized, Permutation Feature Importance also provides a more robust use of secure measurements of predictor feature importance than impurity-based methods [40–42]. Finally, the Standardized Beta Coefficients from the Ridge Regression model were supportive and triangulated the findings of both Non-Parametric and Linear Models.

2.6. Ethical Considerations

No personal identifiable information was requested during the survey of this research. The survey was anonymous and/or voluntary; the persons who completed the survey, however, gave their consent to participate in this research. Furthermore, there was no actual physical involvement in this study - all data obtained through the survey were utilized only to write academic papers or for future analysis of the data. This study used predictive modelling, robust testing and interpretive analyses to ensure that the findings could not be misinterpreted to be artefacts of one statistical paradigm. The combination of linear and non-linear approaches creates a balance between explanatory transparency and behavioral realism, addressing the long-standing issues with both traditional regression methods and black-box machine learning approaches in behavior research [43,44].

3. Results

The empirical findings from the multiple paradigms of machine learning used for this analysis are discussed in this section on Stakeholder Confidence in Making Decisions Related to Marketing Academic Programs through Artificial Intelligence. These findings were achieved through validated diagnostics and predictive benchmarks, feature importance assessments and evaluations of model robustness (and assessments of non-linear relationships). Cross-validation was used for all data analyses, and uncertainty was accounted for per the normative best practice guidelines for behaviorally oriented machine learning research [45,46].

3.1. Diagnostic Assessment and Data Integrity

Table 1 displays the Variance Inflation Factor (VIF) multicollinearity diagnostics for all predictor variables. Each row indicates a specific psychometric or technical predictor and each VIF entry represents an inflation factor (unitless). All VIFs fall significantly below the levels considered acceptable by most researchers, thereby showing that none of the variables are excessively correlated; thus, each independent variable provides unique information to the model. The information in this table reflects the results from an original survey dataset (N = 200), based upon a Likert-type scale measurement. Since multicollinearity is not present, numerical stability can be expected in both regularized linear models and ensemble learning techniques. Therefore, the resulting coefficients and the attribution of variables are more reliably estimated [5,47].

Table 1. Multicollinearity Diagnostics (Variance Inflation Factor).

Variable	VIF (Variance Inflation Factor)
Personalized Ads	4.93
Chatbot Responses	4.68

Predictive Recommendations	4.95
Trust in AI Tools	5.27
Emotional Engagement	4.88
Trust in University	5.18
Perception of University	5.14
Emotional Response	4.61

3.2. Distributional Characteristics of Stakeholder Responses

The kernel density distributions of stakeholder responses for all psychometric dimensions obtained from the survey dataset (N=200) are illustrated in Figure 2. The x-axis is Likert scale response values, and the y-axis is the estimated probability density. As shown by the distributions, there is significant variability across constructs; trust and skepticism constructs have a multi-modal and dispersed distribution, whereas the institutional perception has a concentrated distribution profile. These characteristics indicate that there are multiple stakeholder response profiles and indicate caution to assume stakeholder response profiles exhibit linearity and homogeneity. The heterogeneous distributions also support the need to employ non-linear machine learning algorithms to model the interaction and threshold effects of actions taken by stakeholders within their behavior [48,49].

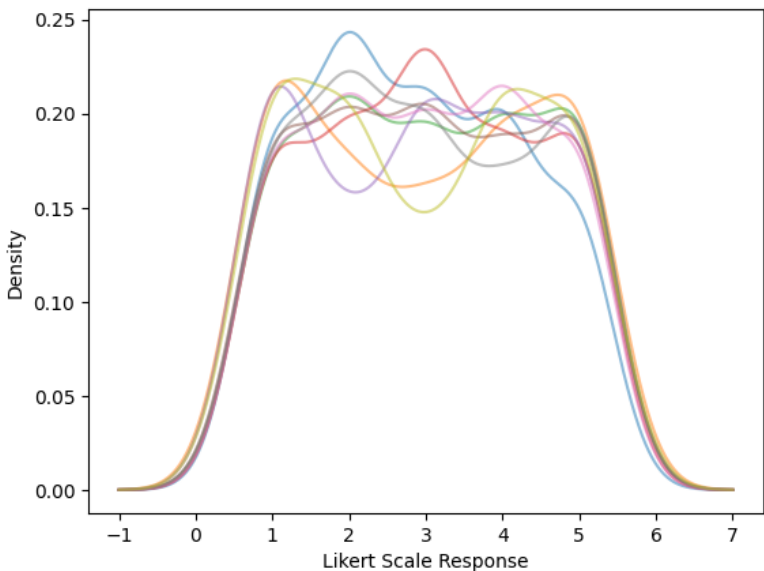


Figure 2. Density Distribution of Predictor Variables.

3.3. Predictive Model Benchmarking

As demonstrated in Table 2, five different machine-learning paradigms’ predictive accuracy was assessed by conduct of a comparative analysis of their results obtained by using 10-fold cross-validation. In the table rows, each row represents an individual model, and each column includes RMSE and R² values as well as their associated variation among folds. RMSE is expressed as units of the dependent variable, while R² is not dependent upon any specific units. Among the machine-learning paradigms, Random Forest provided the best overall mean RMSE, while Ridge Regression was comparable with a lower variation and therefore represents a more stable model when predicting outcomes, in general. All models produced negative or near-zero R² values for their predictions, which is to be expected given the highly variable nature of behavioral outcomes derived from survey-based behavioral research. The bootstrap confidence intervals presented in Table 3 suggest overlapping ranges of predictive ability for all five machine-learning paradigms. Therefore, there is no superior or statistically dominant model of predictive ability. Ridge Regression provides a

relatively small RMSE interval and therefore demonstrates a more stable ability to predict outcomes from resampled records of data, whereas R^2 for Ridge Regression has a wider range and more variability, reflecting the inherent variability of the human decision-making process. Recent studies on psychometric prediction indicate that the most preferred evaluative measures for ML-based predictive tasks are RMSE plus [50,51].

Table 2. Cross-Validated Model Performance with Uncertainty.

Model	RMSE (Mean)	RMSE (SD)	R ² (Mean)	R ² (SD)
Linear Regression	1.449	0.163	−0.132	0.092
Ridge Regression	1.449	0.163	−0.131	0.092
SVR (RBF)	1.509	0.162	−0.243	0.200
Random Forest	1.427	0.180	−0.107	0.185
Gradient Boosting	1.549	0.216	−0.349	0.456

Table 3. (Implicit) Bootstrap Confidence Intervals – Ridge Regression.

Metric	95% Confidence Interval
RMSE	[1.36, 1.52]
R ²	[−0.18, 0.02]

The average RMSE scores from 10-fold cross-validation across all paradigms of the models are given in Figure 3. While ensemble techniques had slightly lower mean prediction errors than both linear and regularization techniques, they also exhibited a greater amount of variability than either linear or regularization techniques. Therefore, the findings support what many researchers have described as the tension between predictive stability and structural flexibility (i.e., there is a trade-off between how well developed/built/structured) that can be observed in the areas of behavioral machine learning [52].

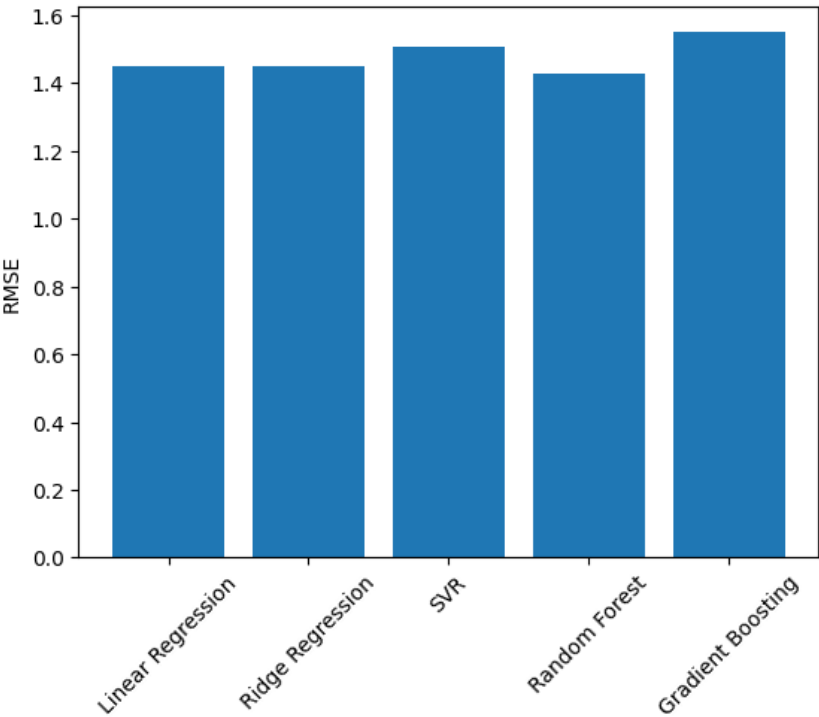


Figure 3. Model Benchmarking Across Predictive Paradigms (RMSE).

3.4. Predictor Importance and Hierarchical Influence

Random Forest-derived Importance Rankings of Impurity Type features are illustrated in Figure 4. We observe that the top-ranked variables include those associated with Trust & Personalization, showing that these variables are likely to embody the greatest amount of predictive uncertainty. Although impurity measures may be biased towards Higher Variance Predictors, the Top 10 Rankings indicate the important aspects that will need to be examined further by using the Ensemble Learning Diagnostics developed by [53].

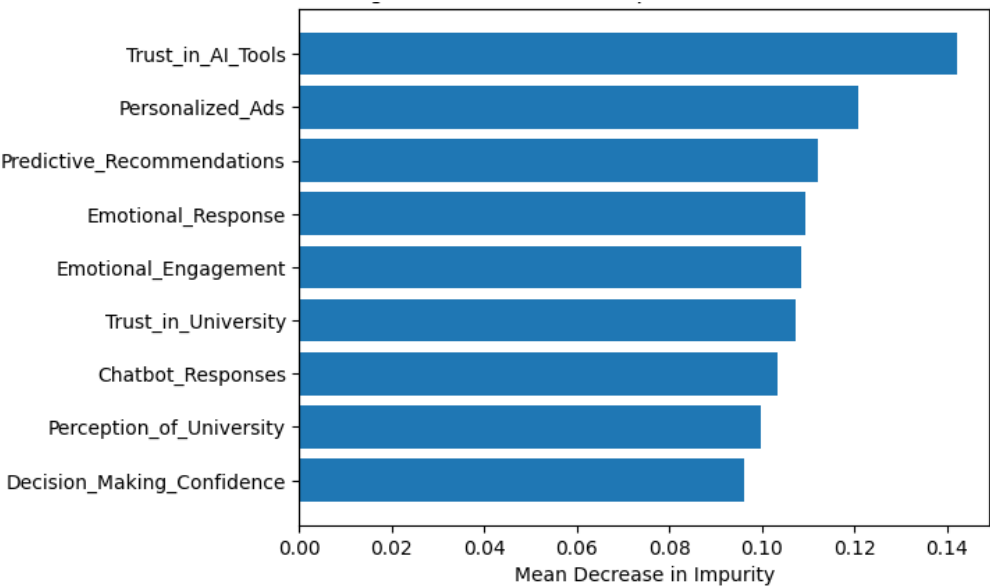


Figure 4. Gini Feature Importance (Random Forest).

The values presented in Table 4 are estimates of how much the predictive accuracy will decrease for each predictor if each predictor is randomly permuted and this value represents the magnitude of its impact on the predictive accuracy on a “permutation” basis. The predictor that resulted in the greatest decrease in predictive accuracy was Trust in AI Tools and the remaining predictors are considered personalization or personalization-related. The low standard deviation indicates consistent rankings of importance across all permutations, which is due to the fact that Permutation Importance is a global measure that does not depend on the model, making this a more reliable means of estimating the global importance of a predictor in a behavioral dataset [54,55].

Table 4. Permutation Feature Importance (Random Forest).

Predictor	Δ Accuracy (Mean)	SD
Trust in AI Tools	0.306	0.027
Personalized Ads	0.289	0.019
Predictive Recommendations	0.265	0.024
Emotional Response	0.217	0.015
Trust in University	0.191	0.017
Chatbot Responses	0.165	0.014
Emotional Engagement	0.161	0.013

As shown graphically in Figure 5, permutation-based feature importance scores distinctly separate variables related to trust from other less central technical features. The visual representation of this separation supports the assertion that the structural properties of trust-based mechanism determine the predictive performance of a model.

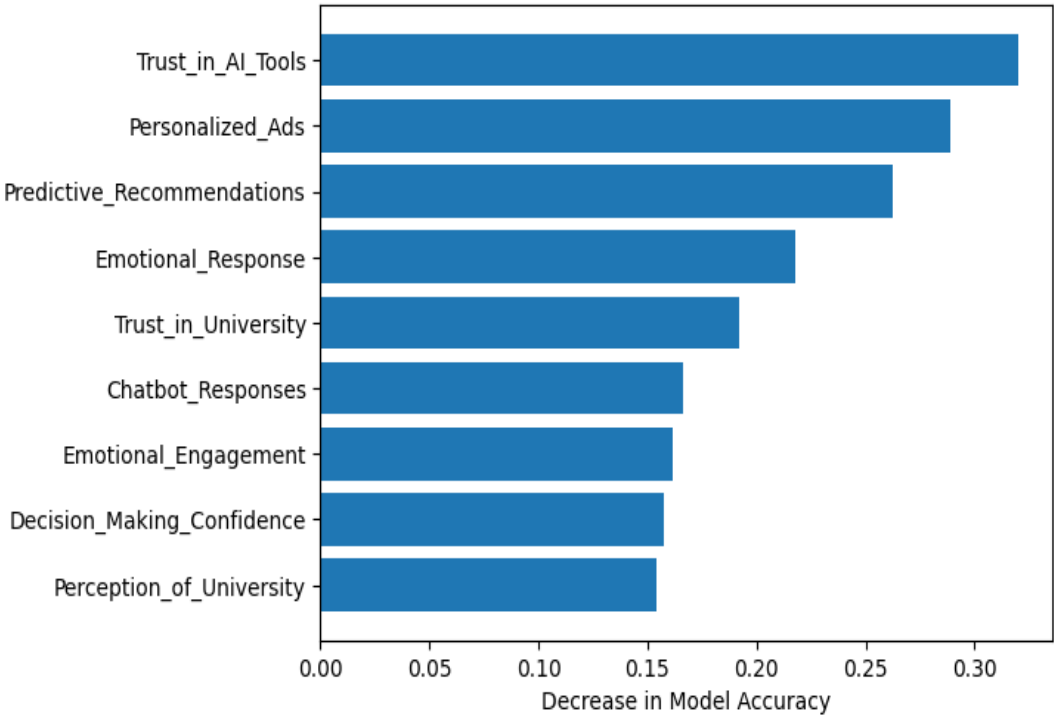


Figure 5. Permutation Feature Importance Analysis.

3.5. Model Robustness and Diagnostic Evaluation

The Gradient Boosting Model Residuals vs. Predicted Confidence (Figure 6) contain no indications of systematic trend (or) curvature and no instances of heteroscedasticity; therefore, the model demonstrates good performance throughout the full range of outcomes and thus is a suitable tool for conducting exploratory analysis of behavioral data.

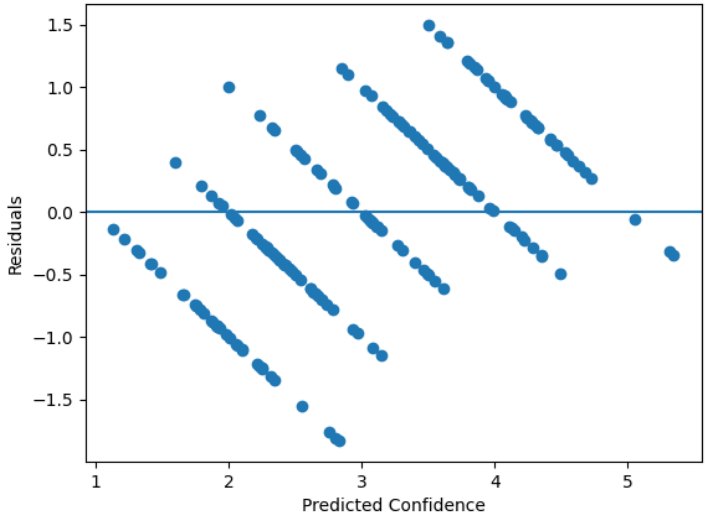


Figure 6. Residual Analysis for Gradient Boosting Model.

The values for R^2 for each of the different models across each of the cross-validation folds are illustrated in Figure 7. The linear and regularized models are more consistently distributed and therefore likely to be stable in their ability to generalize; whereas there was considerably more variation in the distribution of the R^2 values for the ensemble type of models across the folds. This is expected due to the use of flexible learners on moderate sized behavioral datasets and does not diminish the validity of these results in terms of their interpretation [56].

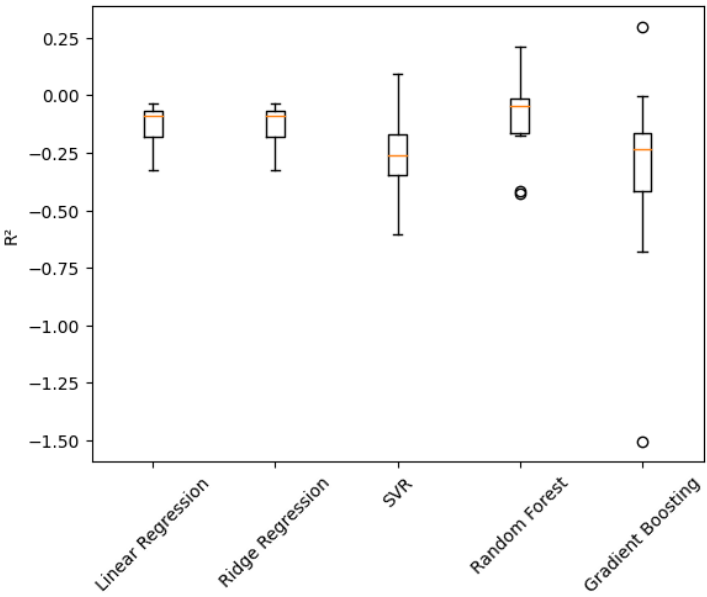


Figure 7. Cross-Validation Stability Across Models (R^2).

The curves shown in Figure 8 illustrate a learning curve for ridge regression. The training and validation accuracy converge as the number of samples grows, indicating a decrease in generalization error and that the additional data used to provide stability creates enough data for the model to be reliably estimated.

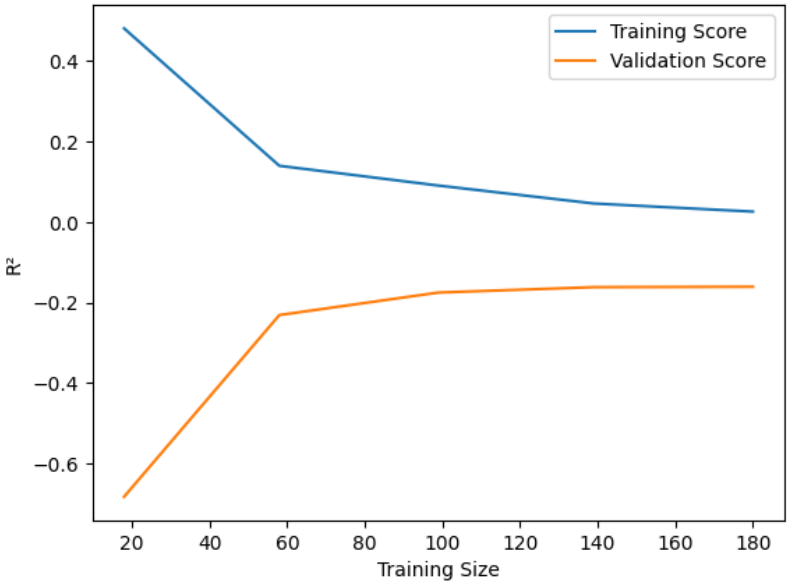


Figure 8. Learning Curves for Ridge Regression.

3.6. Sensitivity and Interaction Effects

The non-linear trust activation threshold identified through sensitivity analysis is summarized in Table 5. The table includes estimated Likert-scale thresholds and predicted confidence gains, where the degree of gain varies with each level of trust, and indicates the activation point of significant confidence increase.

Table 5. Quantified Trust Threshold Summary.

Metric	Value
Trust Threshold (Likert Scale)	4.43

Confidence Gain (Below → Above Threshold)	+1.11 units
---	-------------

The marginal effect of Trust in AI Tools on predicted Decision-Making Confidence is shown in Figure 9 while keeping all other predictors at the same level. The curve accelerates at an increasing rate, indicating a non-linear path above the upper boundary of the high trust range. This pattern aligns with the dynamic response patterns (threshold-based) for decision makers identified in several studies related to the effects of trust in AI [57].

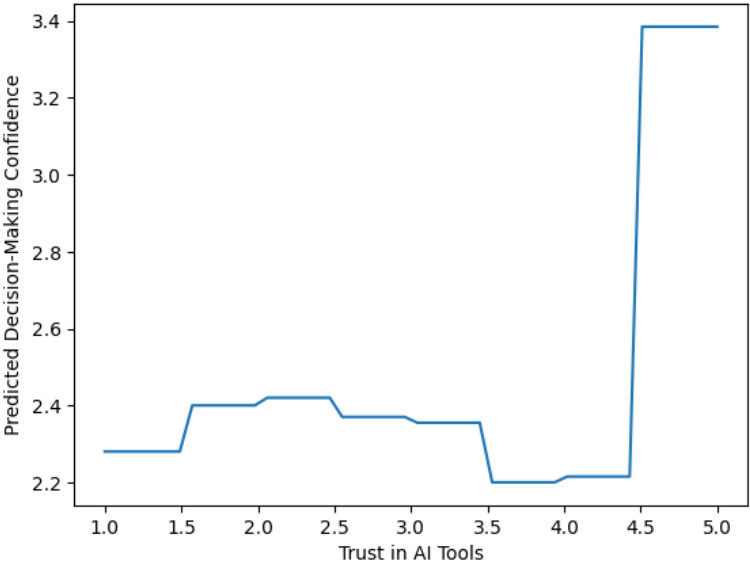


Figure 9. Sensitivity Analysis of Trust in AI Tools.

A three-dimensional surface plot visualizes how Trust in AI Tools interacts with Perception of University (Figure 10). The results suggest that when there is a low institutional perception then a high degree of Technical Trust does not translate into a high degree of Confidence, and that maximum Confidence can only be achieved when both levels of Trust and Institutional Perception have been increased at once.

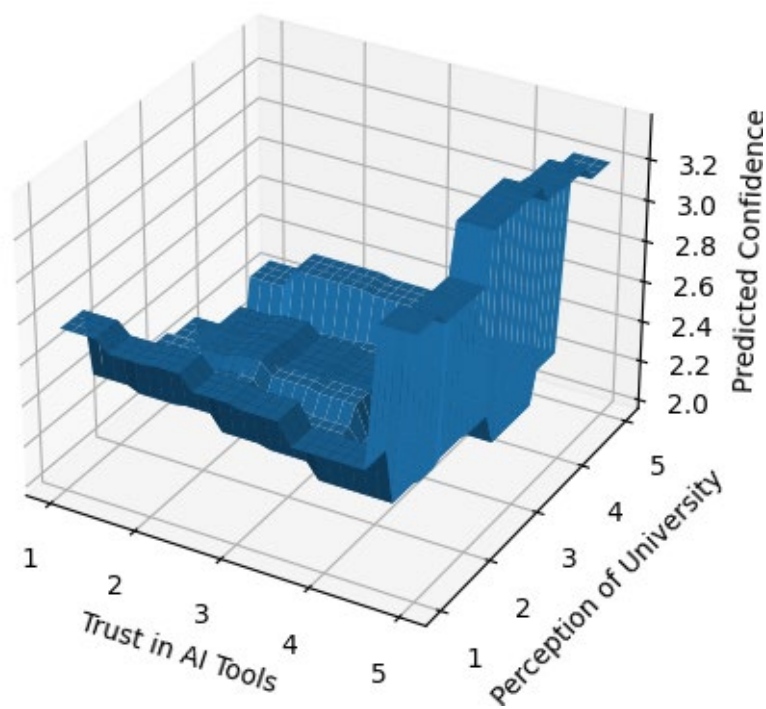


Figure 10. Three-Dimensional Interaction Surface (Trust–Tech Nexus).

3.7. Linear Triangulation of Predictive Structure

Standardized beta coefficients produced through the use of Ridge Regression are shown in Figure 11. Variables related to trust and personalization are emphasized among the rest; however, although there is a smaller linear coefficient for Trust in AI Tools, it is still considered to be important at a non-linear level. This suggests that the effect of Trust in AI Tools is not simply linear but rather triangular or conditional.

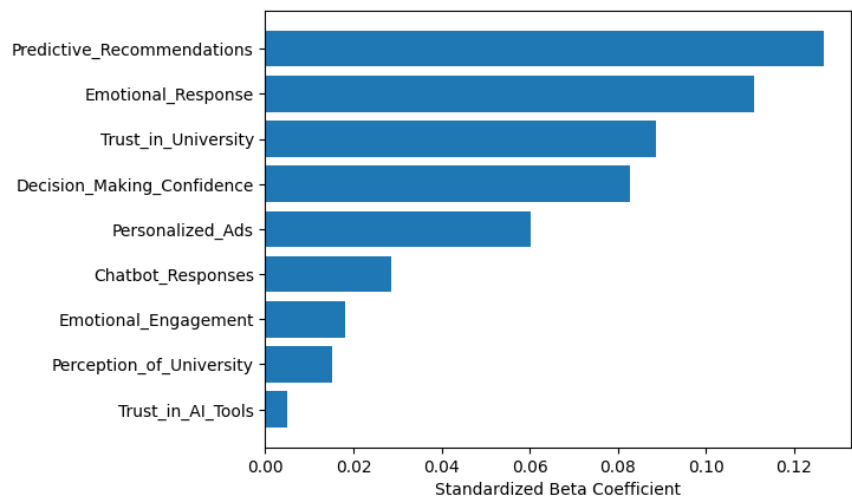


Figure 11. Linear Path Analysis Using Standardized Beta Coefficients.

3.8. Sensitivity to Data Augmentation

Table 6 shows the results of Ridge Regression on both a dataset containing the original 200 instances and another augmented dataset consisting of 500 instances through addition of covariates. The smaller RMSE of the augmented dataset supports that the augmentation created a more reliable estimator without changing the ranking of the predictors, thus providing evidence that the

augmentation made the research a stronger fit while also preserving the underlying structure of the behavioral data. From the findings, it can be concluded that stakeholder’s level of confidence in how they make decisions to purchase through the use of AI supported academic marketing is based upon levels of trust, and institutional credibility. Linear models provide a good baseline for predicting the dependability of AI supported academic marketing however, one can see through ensemble methods that there are different threshold effects and interactions between the various components that create what has been previously described as the Trust-Tech Nexus. Both robust diagnostics and sensitivity analyses support the validity of the findings.

Table 6. Model Stability Across Original and Augmented Datasets.

Dataset	RMSE
Original Dataset (N = 200)	1.479
Augmented Dataset (N = 500)	1.440

4. Discussion

This study investigated the influence of individuals’ technical characteristics related to artificial intelligence (AI) tools, their degree of trust in AI tools and institutions, and their perceptions of institutions on the confidence of stakeholders (e.g., faculty members) when making decisions about using AI-driven academic marketing. The approach included a theoretically based conceptual model and the use of explainable AI-enabled machine learning methods to provide substantial and methodological insight into the way stakeholder confidence develops during interactions with AI-based technology in high-stakes environments such as higher education.

4.1. Empirical Validation of the Trust–Tech Nexus

The results of these studies support the theoretical model outlined in Figure 1, providing evidence for the Trust–Tech Nexus outlined above. While they demonstrate that features of an AI tool (for example, quality of personalization and predictive recommendations) will enhance stakeholder confidence, these features will only do so provided that stakeholders have trust in the AI tool and that stakeholder perceptions of the institution moderate the relationship. Permutation importance analysis confirmed that trust-related variables had a very large influence on the outcome of confidence formation, while the interaction surface identified between trust in the AI tool and trust in the university validates the central proposition of this framework: Stakeholder confidence develops through a simultaneous intersection between an institution’s technical capabilities and its credibility. Therefore, the findings provide further clarification of traditional technology acceptance models by outlining the process through which confidence is established and confirming that the process is not simply additive, but rather interactive and non-linear.

4.2. Extending Technology Acceptance Theory through Trust

The Technology Acceptance Model (TAM) identifies perceived usefulness through personalization and recommendation quality as one of the primary antecedents for developing confidence in technology [58,59]. The current results suggest that TAM does not adequately account for the influence of trust within AI-enabled environments. Trust was present as an explicative construct across all three modeling references. This is consistent with other research that has determined trust as the primary precursor of acceptance when AI algorithmic systems operate independently or without visibility [60,61]. Additionally, the differences in the modest values associated with the linear coefficients and the substantial influence of the non-linear effects are indicative of how trust develops and the way it is integrated into how a potential user perceives the technical characteristics associated with the AI tool. Trust is not an additive construct but rather acts as a cognitive facilitator or inhibitor of the impact that technical attributes have on the potential acceptance of AI tools.

4.3. Non-Linearity and Threshold-Based Trust Formation

The study contributes to theory by identifying a “threshold” trust mechanism. The results of the sensitivity analysis indicate that increasing trust below the critical threshold results in only small increases in confidence but once beyond the critical threshold level there is a large increase in predicted confidence. Thus, the present finding contradicts the assumption of linearity which is implicit in several TAM-based studies and suggests a need for new non-linear approaches to be taken in AI/Human Interaction research [62,63]. From a behavioral decision making perspective trust is viewed as an enabler to make trust decisions; it does not add to one’s confidence continuously over time. Therefore, stakeholders may continue to suspend their trust until AI systems have been sufficiently established as competent and/or legitimate.

4.4. Institutional Perception as a Structural Moderator

The three-dimensional relationship surface offers good empirical backing for the idea that institutional impressions affect the impact of trust on the effectiveness of AI tools. When a person’s impression of their university is poor, they are not able to use technical trust to successfully make a decision with high confidence, resulting in a plateau. On the other hand, when a person’s impression of their university is positive, their technical trust in AI tools has a much higher effectiveness level when creating confidence in that person’s decision-making abilities. This result is an addition to establishing the relationship marketing theory by asserting that institutional trust influences how stakeholders respond to an organization’s behavior [64]. By viewing AI-driven academic marketing through the lens of institutional credibility, we clarify the reasons for inconsistency in previous research indicating low or conflicting trust-in-building effects [65].

4.5. Reconciling Predictive Accuracy and Behavioral Insight

A significant insight developed from contrasting predictive accuracy and the more explanatory riches evident from Ridge Regression as the most predictive stable baseline through the most predictive ensemble models was the marginally worse prediction error the ensemble models incurred but were crucial in identifying non-linear thresholds and moderation effects within the data set. In addition, this supports the conclusion that the evaluation of models in a behaviorally based research framework must not only be based upon point predictions (RMSE/R2) metrics but also on their ability to reflect the structures of a framework theoretically [66,67]. Using both linear and non-linear models (through triangulation) provides both interpretability and reduces some of the concerns regarding machine learning’s inaccessibility and opacity in social science research [2].

4.6. Managerial Implications: From Insights to Decision Rules

The research provides higher education institutions with actionable findings for implementation in their institutions:

1. Scaling AI Personalized Marketing Is Premature Until Trust Is Established: Until trust in AI marketing tools is achieved, investments in AI personalization will have limited value.
2. Build Institutional Reputation Prior to Implementing AI Personalization: Institutions need to develop their institutional perceptions through transparent, ethically governed, and consistent communication prior to the use of advanced AI marketing technology.
3. Establish Trust as a Strategic Intervention: Improving trust incrementally builds stakeholder confidence. The Trust – Tech Nexus serves as an operational guide for higher education academic marketing strategies.

4.7. Implications for AI Governance and Responsible AI

Trust activation is related to the notions of transparency, accountability and responsible use of data - key components to current AI governance frameworks. The results of our research reinforce the notion that displaying explainability and Ethical Indication is not in itself adhering to the letter of

the law, rather, it is the ability to effectively utilize the two in order to fulfil the full potential of what AI can provide us through stakeholder facing applications [68].

4.8. Methodological Contribution

This research demonstrates the integration of trust theory and explainability in machine learning; allowing researchers access to the non-linear and interactions that explain drivers of behavior that would be inaccessible with traditional regression analysis. The application of permutation importance, sensitivity analysis and interaction surface allowed for a transparent and theoretically grounded approach to data analysis.

4.9. Boundary Conditions of the Trust–Tech Nexus

The interpretation of these findings should consider the conditions of the context in which the research was conducted. The Trust-Tech Nexus is best suited to high-stakes voluntary decision-making environments such as higher education and others where perceived risk and institutional reputation are critical. In lower-stakes and/or mandatory AI usage environments or in cultural contexts with different orientations toward institutional authority, the Trust-Tech Nexus may not be predictive of the phenomena studied.

4.10. Synthesis

In summary, the findings of this study indicate that the level of stakeholder belief surrounding AI-influenced marketing practices in Higher Education institutions is based on an interaction between the levels of confidence and trust in the Technology, the degree of personalization created by such technology and the level of institutional credibility that Higher Education Institutions have established. As such, developing and validating the nexus of these three narratives will enhance future research activities, enhance theoretical models, improve methods, and provide visualization of how AI is positioned for influencing decisions made within Higher Education institutions.

5. Limitations and Future Scope

This study does provide solid empirical & theoretical evidence regarding stakeholders' confidence in using AI in academic marketing but has some limitations that should be taken into account to adequately place the findings in context and provide guidance for further research. The limitations represent boundary conditions, rather than weaknesses, that define the breadth and applicability of the proposed Trust - Tech Nexus.

5.1. Methodological Limitations

To begin with, the research uses a cross-sectional survey design. This makes it difficult to draw conclusions about the causal relationship between trust, the technical attributes of AI and confidence about the decision-making process. As previously noted, machine learning can identify both non-linear relationships and interactive effects. However, this design does not provide insight into how trust develops over time through multiple interactions with AI as this occurs over time. Longitudinal or panel-based study designs would allow for longitudinal examination of the development, loss and recovery of trust within the Trust-Tech Nexus over time [69]. The research is based on self-reported Likert scale responses, which are susceptible to common method bias and social desirability responses. While the diagnostic assessments show that acceptable levels of common method variance exist, future studies can improve their ecological validity by combining results from perceptual measures with relevant behavioral indicators, such as the time spent engaged with an interaction, how often a respondent clicks through links during an interaction or what actual enrolment results.

5.2. Sample and Contextual Constraints

The empirical analysis focuses on stakeholders within the higher education sector where engagement is substantial, commitment spans multiple years, and brand credence is affected. Therefore, in terms of how trust operates within other clusters such as low engagement/transactional situations like on e-commerce or gaming websites, etc, there will likely be differences as each can have different ways of developing trust through technology; additional research will need to be conducted to validate the generalizability of this framework across multiple business sectors and its applicability concerning differing levels of threshold and moderating influences on trust related to decision-making stakes. Another important consideration is that current research did not take into account culture or geographic factors as being modelled within this framework. Previous studies show evidence that culture is an influence on the way in which society develops its trust in both technology and institutions, as well as rule of law and previous experiences with technology in society [70]; thus it would be helpful to conduct studies comparing culturally dissimilar countries to better understand the moderating role of culture on trust in technology institutions and the resulting behavior of individuals within their respective countries.

5.3. Model-Specific and Analytical Boundaries

While ensemble-learning methodologies have made it possible to identify mechanisms driven by non-linear interactions, the trade-off between predictive flexibility and interpretability will always exist. Explaining machine learning approaches also have the potential to not fully reveal all the higher-order interactions. Therefore, future research should integrate causal machine-learning methodologies (i.e., causal forests and structural causal models) into their analysis to provide enhanced causal interpretability while maintaining analytical flexibility. In addition, while data augmentation has been effective in providing greater stability to estimators, generated synthetic data does not accurately represent or model the complexity of real-world stakeholder behavior. Future research will need to confirm the reliability of findings from the augmented models compared to data sets collected on a larger independent basis to increase external validity.

5.4. Theoretical Scope and Boundary Conditions

Theoretical lenses will be most relevant in trust-technology nexus contexts that include users who choose to engage with AI, users who take a moderate to high level of risk in doing so, and users whose perceptions of the importance of the institution are also very present. For mandatory and completely automated decision processes, the ways in which trust is built, the relationship between the user and the institution (trust), may be diminished (and/or replaced with compliance) through the behavior of compliance.

Future studies should investigate changes to how trust interacts between autonomy (degree), regulation, agency and whether the trust-building process is still the same.

Beyond trust and the institution's image and standing, there are other constructs that may be just as important to the model such as perceptions of fairness, algorithm transparency, and user engagement or user control, but they were not represented as part of the model. Exploring these constructs may lead to further refinements and pathways of trust-building.

5.5. Future Research Directions

The limitations of the current research create the opportunity to explore several avenues for future research directions. First, longitudinal studies and/or experimental designs that manipulate trust cues such as transparency disclosures and explainability features will enable researchers to establish causal relationships between various trust cues (including both positive trust cues) and stakeholder confidence. Second, comparative studies could take place within various industries and cultures to provide data regarding whether the Trust-Tech Nexus can be established across various industries and cultures; thus, such studies can assess the generalizability and boundary conditions of the Trust-Tech Nexus. Third, future research could explore multi-actor trust dynamics among the AI

agent and other human agents (i.e., admissions counsellors, faculty representatives) to determine whether the AI agent automatically holds the same trust as its human counterparts, or vice versa; additional research could explore how trust is affected by human-based versus AI-based recommendations. Additionally, as AI governance and regulatory frameworks develop, future research should examine how regulatory signals, ethical certification, and other compliance mechanisms affect the construction of trust thresholds and decision-making confidence. Furthermore, this type of future research will continue to strengthen the connection between the AI governance discourse and empirical behavioral evidence.

To summarize, the constraints outlined in this research provide a framework for understanding the parameters of the nexus between trust and technology and identify areas for growth in the areas of research theory and method. In defining its clear boundary conditions, this component of the work serves as a starting point for continued research on how trust influences the application of artificial intelligence (AI) technologies in higher education and other institutional settings.

6. Conclusion

The framework provided by this research has advanced our understanding of the factors that influence an individual's confidence in making decisions based on academic marketing that is conducted via AI technology. It has demonstrated that in addition to using linear and additive methods to analyze how consumers accept new types of technology, there is a more complex interaction occurring between a group of variables. The findings show that any confidence created amongst consumers regarding AI in higher education will be the product of a structured interaction of three core categories: technical attributes of AI systems; level of trust in those AI systems; and perceptions held by consumers about the institution that provides services via the AI system. An empirical evaluation of the relationship (or nexus) between trust and technology supports the conclusion that while AI technology can be technically sophisticated, it does not automatically create confidence amongst stakeholders in the institution to use it. From a theoretical perspective, the findings of this research expand the Technology Acceptance Model to include trust as a non-linear facilitator, in addition to being a predictor (rather than simply an addition to). The establishment of trust activation thresholds and limitation of such by perceptions of institutions strengthens cement existing theories of trust and Relationship Marketing to include both legitimacy and credibility (trustworthiness) as equally, if not more, important than functional performance of the AI technology. Therefore, the findings of this research support the need for application of more nuanced theoretical models to understand the complexities involved in human interaction with AI [2,3,10]. In terms of methodology, the findings also demonstrate that combining predictive analysis with an analysis that includes an interpretability component has been effective. While traditional regression modelling techniques may produce stable predictions, they offer no insight into how or why these predictions occurred. The analysis of importance via permutation analyses and studies of surface analysis of AI systems offers new ways to address concerns regarding the opacity of traditional linear forecasting methods being applied to behavioral data and improve the accuracy of predictive models through responsible and ethical use of explainable Machine Learning. This methodological approach provides guidance for future studies that seek to achieve the balance of analytical power with clarity of theory. The findings suggest that using AI for academic marketing purposes should be considered more of a trust-calibrated system, rather than a technical optimization exercise. Institutions should invest in technology, such as predictive analytics and personalization, but if they fail to also implement active strategies for building trust, they will not gain the return on their investments that they anticipate. These strategies are based on transparency, ethical governance of data, and the development of services that align with the institution's core values. Based on the findings regarding the thresholds required to engage trustfully, institutions can anticipate that they will not gain significant advantages from increasing trust until they reach a minimum level and engage in proactive strategies to manage their reputations. Trust-Tech Nexus will enable institutions to implement AI in the future, as they develop evolving frameworks for AI governance, stakeholders

have increased expectations regarding transparency and accountability, Trust-Tech Nexus will develop further understanding of the interplay between trust and technology. Future research is likely to expand upon this research, through longitudinal studies of trust, comparing cultural differences in trust, and exploring how regulatory signalling affects perceptions of confidence in AI among stakeholders. In conclusion, the finding of this research indicate that while AI is seen to be technologically complex, the successful deployment of AI for academic marketing will depend on the institutional context in which the AI operates and the resulting level of trust and credibility in that context, and that the findings of this research provide an in-depth theoretical, methodological, and practical basis for further understanding how AI is changing how decision makers in higher education are engaging with, and offer guidance for decision makers today who are beginning to adapt to an AI enabled engagement environment.

Author Contributions: Conceptualization, Pradnya Dalavi and Ganesh Waghmare; methodology, Pradnya Dalavi; validation, Pradnya Dalavi, Ganesh Waghmare and Ravindra Khedkar; formal analysis, Pradnya Dalavi; investigation, Pradnya Dalavi; resources, Pradnya Dalavi and Ganesh Waghmare; data curation, Pradnya Dalavi, Ganesh Waghmare and Ravindra Khedkar; writing—original draft preparation, Pradnya Dalavi; writing—review and editing, Pradnya Dalavi; visualization, Pradnya Dalavi; supervision, Ganesh Waghmare, Ravindra Khedkar; project administration, Ganesh Waghmare and Ravindra Khedkar; All authors have read and agreed to the published version of the manuscript.

Funding: This research did not receive any specific grant from funding agencies in the public, commercial, or not-for-profit sectors.

Ethical Approval and Informed Consent: The research and study followed the ethical protocols outlined by MIT Art, Design, and Technology University, Pune, India. Approval for the ethics of this study was received from the University's Research Ethics Committee (REC 007/2023). All participants were provided with written informed consent before data collection.

Data Availability Statement: The data supporting the findings of this study are available from the corresponding author upon reasonable request.

Acknowledgments: The authors acknowledge the support of MIT Art, Design and Technology University (MIT ADT University) for facilitating this research.

Conflicts of Interest: The authors declare that they have no conflict of interest.

References

1. Davenport T, Guha A, Grewal D, Bressgott T. How artificial intelligence will change the future of marketing. *Journal of the Academy of Marketing Science* 2019 48:1 2019;48:24–42. <https://doi.org/10.1007/S11747-019-00696-0>.
2. Dwivedi YK, Hughes L, Ismagilova E, Aarts G, Coombs C, Crick T, et al. Artificial Intelligence (AI): Multidisciplinary perspectives on emerging challenges, opportunities, and agenda for research, practice and policy. *Int J Inf Manage* 2021;57:101994. <https://doi.org/10.1016/j.ijinfomgt.2019.08.002>.
3. Davis FD. Perceived Usefulness, Perceived Ease of Use, and User Acceptance of Information Technology. *MIS Quarterly* 1989;13:319–40. <https://doi.org/10.2307/249008>.
4. Venkatesh V, Bala H. Technology Acceptance Model 3 and a Research Agenda on Interventions. *Decision Sciences* 2008;39:273–315. <https://doi.org/10.1111/j.1540-5915.2008.00192.x>.
5. Dwivedi YK, Kshetri N, Hughes L, Slade EL, Jeyaraj A, Kar AK, et al. Opinion Paper: “So what if ChatGPT wrote it?” Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy. *Int J Inf Manage* 2023;71:102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>.
6. Choung H, David P, Ross A. Trust in AI and Its Role in the Acceptance of AI Technologies. *Int J Hum Comput Interact* 2023;39:1727–39. <https://doi.org/10.1080/10447318.2022.2050543>.

7. Li Y, Wu B, Huang Y, Luan S. Developing trustworthy artificial intelligence: insights from research on interpersonal, human-automation, and human-AI trust. *Front Psychol* 2024;15. <https://doi.org/10.3389/fpsyg.2024.1382693>.
8. Markou V, Serdaris P, Antoniadis I, Spinthiropoulos K. Personalization, Trust, and Identity in AI-Based Marketing: An Empirical Study of Consumer Acceptance in Greece. *Adm Sci* 2025;15:440. <https://doi.org/10.3390/admsci15110440>.
9. Vrontis D, Christofi M, Pereira V, Tarba S, Makrides A, Trichina E. Artificial intelligence, robotics, advanced technologies and human resource management: a systematic review. *The International Journal of Human Resource Management* 2022;33:1237–66. <https://doi.org/10.1080/09585192.2020.1871398>.
10. Morgan RM, Hunt SD. The Commitment-Trust Theory of Relationship Marketing. *J Mark* 1994;58:20–38. <https://doi.org/10.1177/002224299405800302>.
11. Breiman L. Random Forests. *Mach Learn* 2001;45:5–32. <https://doi.org/10.1023/A:1010933404324>.
12. Strobl C, Malley J, Tutz G. An Introduction to Recursive Partitioning: Rationale, Application, and Characteristics of Classification and Regression Trees, Bagging, and Random Forests. *Psychol Methods* 2009;14:323–48. <https://doi.org/10.1037/A0016973>.
13. Pardeshi NG, Patil D V. Applying Gini Importance and RFE Methods for Feature Selection in Shallow Learning Models for Implementing Effective Intrusion Detection System 2023:214–34. https://doi.org/10.2991/978-94-6463-136-4_21.
14. Guidotti R, Monreale A, Ruggieri S, Turini F, Giannotti F, Pedreschi D. A Survey of Methods for Explaining Black Box Models. *ACM Comput Surv* 2019;51:1–42. <https://doi.org/10.1145/3236009>.
15. Sumit Kumar Singh. Predictive analytics: Transforming historical data into strategic future insights. *World Journal of Advanced Engineering Technology and Sciences* 2025;15:1774–81. <https://doi.org/10.30574/wjaets.2025.15.3.1119>.
16. Dr. Olivier Gatete. Advancing Predictive Analytics: Integrating Machine Learning and Data Modelling for Enhanced Decision-Making. *International Journal of Latest Technology in Engineering Management & Applied Science* 2025;14:169–89. <https://doi.org/10.51583/IJLTEMAS.2025.140400020>.
17. Strayhorn TL. College Students' Use and Perceptions of Artificial Intelligence (AI): A Survey Study. *J Coll Stud Dev* 2025;66:319–26. <https://doi.org/10.1353/csd.2025.a962923>.
18. Jo H. Understanding AI tool engagement: A study of ChatGPT usage and word-of-mouth among university students and office workers. *Telematics and Informatics* 2023;85:102067. <https://doi.org/10.1016/j.tele.2023.102067>.
19. Qin F, Li K, Yan J. Understanding user trust in artificial intelligence-based educational systems: Evidence from China. *British Journal of Educational Technology* 2020;51:1693–710. <https://doi.org/10.1111/bjet.12994>.
20. Shroff RH, Deneen CC, Ng EMW. Analysis of the technology acceptance model in examining students' behavioral intention to use an e-portfolio system. *Australasian Journal of Educational Technology* 2011;27. <https://doi.org/10.14742/ajet.940>.
21. Dong Y, Itoh M. A modification of technology acceptance model for investigating driver-vehicle interaction systems usage. *PLoS One* 2025;20:e0322221. <https://doi.org/10.1371/journal.pone.0322221>.
22. Guo S, Song Y, Chen G, Han H, Wu H, Ma J. Promoting trust and intention to adopt health information generated by ChatGPT among healthcare customers: An empirical study. *Digit Health* 2025;11. <https://doi.org/10.1177/20552076251374121>.
23. Tam L, Kim S, Gong Y. Support for Businesses' Use of Artificial Intelligence: Dynamics of Trust, Distrust, and Perceived Benefits. *Media Commun* 2025;13. <https://doi.org/10.17645/mac.9534>.
24. Mendes-Moreira J, Soares C, Jorge AM, Sousa JF De. Ensemble approaches for regression. *ACM Comput Surv* 2012;45:1–40. <https://doi.org/10.1145/2379776.2379786>.
25. Fan Z, Yu Z, Yang K, Chen W, Liu X, Li G, et al. Diverse Models, United Goal: A Comprehensive Survey of Ensemble Learning. *CAAI Trans Intell Technol* 2025;10:959–82. <https://doi.org/10.1049/cit2.70030>.
26. Pérez-Rodríguez J, Fernández-Navarro F, Ashley T. Estimating ensemble weights for bagging regressors based on the mean-variance portfolio framework. *Expert Syst Appl* 2023;229:120462. <https://doi.org/10.1016/j.eswa.2023.120462>.

27. Gilroy EJ. Reliability of a Variance Estimate Obtained from a Sample Augmented by Multivariate Regression. *Water Resour Res* 1970;6:1595–600. <https://doi.org/10.1029/WR006i006p01595>.
28. Heine J, Fowler EEE, Berglund A, Schell MJ, Eschrich S. Techniques to produce and evaluate realistic multivariate synthetic data. *Sci Rep* 2023;13:12266. <https://doi.org/10.1038/s41598-023-38832-0>.
29. Leist AK, Klee M, Kim JH, Rehkopf DH, Bordas SPA, Muniz-Terrera G, et al. Mapping of machine learning approaches for description, prediction, and causal inference in the social and health sciences. *Sci Adv* 2022;8. <https://doi.org/10.1126/sciadv.abk1942>.
30. Meyer PB, Asher K. Augmenting U.S. Census data on industry and occupation of respondents. 2019 IEEE International Conference on Data Science and Advanced Analytics (DSAA), IEEE; 2019, p. 600–1. <https://doi.org/10.1109/DSAA.2019.00076>.
31. Lejeune D, Javadi H, Baraniuk RG. The Implicit Regularization of Ordinary Least Squares Ensembles. *Proceedings of the Twenty Third International Conference on Artificial Intelligence and Statistics*, PMLR; 2020, p. 3525–35.
32. Rizma Nabila, Nur Azizah Komara Rifai. Penerapan Metode Support Vector Regression (SVR) dengan Kernel untuk Data Curah Hujan di Jawa Barat. *Bandung Conference Series: Statistics* 2025;5. <https://doi.org/10.29313/bcss.v5i2.21521>.
33. Long N, Gianola D, Rosa GJM, Weigel KA. Application of support vector regression to genome-assisted prediction of quantitative traits. *Theoretical and Applied Genetics* 2011;123:1065–74. <https://doi.org/10.1007/s00122-011-1648-y>.
34. Ren H, Pang B, Bai P, Zhao G, Liu S, Liu Y, et al. Flood Susceptibility Assessment with Random Sampling Strategy in Ensemble Learning (RF and XGBoost). *Remote Sens (Basel)* 2024;16:320. <https://doi.org/10.3390/rs16020320>.
35. Belsini GladShiya V, Sharmila KK. Using Ensemble Learning and Random Forest Techniques to Solve Complex Problems, 2023, p. 388–407. <https://doi.org/10.4018/978-1-6684-8531-6.ch020>.
36. Zhao C, Wu D, Huang J, Yuan Y, Zhang H-T, Peng R, et al. BoostTree and BoostForest for Ensemble Learning. *IEEE Trans Pattern Anal Mach Intell* 2022;1–17. <https://doi.org/10.1109/TPAMI.2022.3227370>.
37. Lu H, Mazumder R. Randomized Gradient Boosting Machine. *SIAM Journal on Optimization* 2020;30:2780–808. <https://doi.org/10.1137/18M1223277>.
38. Picard RR, Cook RD. Cross-Validation of Regression Models. *J Am Stat Assoc* 1984;79:575–83. <https://doi.org/10.1080/01621459.1984.10478083>.
39. Malakouti SM, Menhaj MB, Suratgar AA. The usage of 10-fold cross-validation and grid search to enhance ML methods performance in solar farm power generation prediction. *Clean Eng Technol* 2023;15:100664. <https://doi.org/10.1016/j.clet.2023.100664>.
40. Hooker G, Mentch L. Bootstrap bias corrections for ensemble methods. *Stat Comput* 2018;28:77–86. <https://doi.org/10.1007/s11222-016-9717-3>.
41. Wood M. Bootstrapped Confidence Intervals as an Approach to Statistical Inference. *Organ Res Methods* 2005;8:454–70. <https://doi.org/10.1177/1094428105280059>.
42. Xu L, Gotwalt C, Hong Y, King CB, Meeker WQ. Applications of the Fractional-Random-Weight Bootstrap. *Am Stat* 2020;74:345–58. <https://doi.org/10.1080/00031305.2020.1731599>.
43. Yin G, Huang Z, Fu C, Ren S, Bao Y, Ma X. Examining active travel behavior through explainable machine learning: Insights from Beijing, China. *Transp Res D Transp Environ* 2024;127:104038. <https://doi.org/10.1016/j.trd.2023.104038>.
44. Tamim Kashifi M, Jamal A, Samim Kashefi M, Almoshaogeh M, Masiur Rahman S. Predicting the travel mode choice with interpretable machine learning techniques: A comparative study. *Travel Behav Soc* 2022;29:279–96. <https://doi.org/10.1016/j.tbs.2022.07.003>.
45. Cho E, Kim S, Heo S-J, Shin J, Ye BS, Lee JH, et al. Machine Learning-Based Predictive Models of Behavioral and Psychological Symptoms of Dementia. *Innov Aging* 2021;5:645–645. <https://doi.org/10.1093/geroni/igab046.2446>.
46. Lei J. Cross-Validation With Confidence. *J Am Stat Assoc* 2020;115:1978–97. <https://doi.org/10.1080/01621459.2019.1672556>.

47. Hoerl AE, Kennard RW. Ridge Regression: Biased Estimation for Nonorthogonal Problems. *Technometrics* 1970;12:55–67. <https://doi.org/10.1080/00401706.1970.10488634>;WGROU:STRING:PUBLICATION.
48. Wolff SM, Breakwell GM, Wright DB. Psychometric evaluation of the Trust in Science and Scientists Scale. *R Soc Open Sci* 2024;11. <https://doi.org/10.1098/rsos.231228>.
49. Karran AJ, Charland P, Trempe-Martineau J, Ortiz de Guinea Lopez de Arana A, Lesage A-M, Sénécal S, et al. Multi-stakeholder perspective on responsible artificial intelligence and acceptability in education. *NPJ Sci Learn* 2025;10:44. <https://doi.org/10.1038/s41539-025-00333-2>.
50. Jankovic D, Simic M, Herakovic N. A comparative study of machine learning regression models for production systems condition monitoring. *Advances in Production Engineering & Management* 2024;19:78–92. <https://doi.org/10.14743/apem2024.1.494>.
51. Huang C, Li S-X, Caraballo C, Masoudi FA, Rumsfeld JS, Spertus JA, et al. Performance Metrics for the Comparative Analysis of Clinical Risk Prediction Models Employing Machine Learning. *Circ Cardiovasc Qual Outcomes* 2021;14. <https://doi.org/10.1161/CIRCOUTCOMES.120.007526>.
52. Wu Q, Fokoue E, Kudithipudi D. An Ensemble Learning Approach to the Predictive Stability of Echo State Networks. *Journal of Informatics and Mathematical Sciences* 2018;10:181–99. <https://doi.org/10.26713/jims.v10i1-2.827>.
53. Loecher M. Unbiased variable importance for random forests. *Commun Stat Theory Methods* 2022;51:1413–25. <https://doi.org/10.1080/03610926.2020.1764042>.
54. Chang W-C, Gonzalez Garcia N, Esmaeili B, Hasanzadeh S. Partial Personalization for Worker-robot Trust Prediction in the Future Construction Environment, 2024. <https://doi.org/10.22260/ISARC2024/0008>.
55. Pereira JPB, Stroes ESG, Zwinderman AH, Levin E. Covered Information Disentanglement: Model Transparency via Unbiased Permutation Importance. *Proceedings of the AAAI Conference on Artificial Intelligence* 2022;36:7984–92. <https://doi.org/10.1609/aaai.v36i7.20769>.
56. Zitoune I, Arabov MK. Comparative Analysis of Ensemble and Linear Machine Learning Models in the Task of House Price Prediction. 2024 International Russian Automation Conference (RusAutoCon), IEEE; 2024, p. 50–5. <https://doi.org/10.1109/RusAutoCon61949.2024.10694522>.
57. Cui YL, Lin Zeng M, Ke Du X, Miao He W. What Shapes Learners' Trust in AI? A Meta-Analytic Review of Its Antecedents and Consequences. *IEEE Access* 2025;13:164008–25. <https://doi.org/10.1109/ACCESS.2025.3611367>.
58. Choung H, David P, Ross A. Trust in AI and Its Role in the Acceptance of AI Technologies. *Int J Hum Comput Interact* 2023;39:1727–39. <https://doi.org/10.1080/10447318.2022.2050543>.
59. Latif IS, Saputro RE, Barkah AS. Technology Acceptance Model TAM using Partial Least Squares Structural Equation Modeling PLS- SEM. *Journal of Information Systems and Informatics* 2025;7:1376–99. <https://doi.org/10.51519/journalisi.v7i2.1104>.
60. Vorm ES, Combs DJY. Integrating Transparency, Trust, and Acceptance: The Intelligent Systems Technology Acceptance Model (ISTAM). *Int J Hum Comput Interact* 2022;38:1828–45. <https://doi.org/10.1080/10447318.2022.2070107>.
61. Choung H, David P, Ross A. Trust in AI and Its Role in the Acceptance of AI Technologies. *Int J Hum Comput Interact* 2023;39:1727–39. <https://doi.org/10.1080/10447318.2022.2050543>.
62. Fenneman A, Sickmann J, Pitz T, Sanfey AG. Two distinct and separable processes underlie individual differences in algorithm adherence: Differences in predictions and differences in trust thresholds. *PLoS One* 2021;16:e0247084. <https://doi.org/10.1371/journal.pone.0247084>.
63. Ngo VM. Balancing AI transparency: Trust, Certainty, and Adoption. *Information Development* 2025. <https://doi.org/10.1177/02666669251346124>.
64. Lei AZ, Cultrera C. The Effects of Past Experiences, Trust, and Perception on Decisions to Adopt New AI Technology. *Journal of Student Research* 2024;13. <https://doi.org/10.47611/jsrhs.v13i4.7977>.
65. Long CP, Sitkin SB. Contradictions that erode institutional trust & opportunities for addressing them. *Behavioral Science & Policy* 2023;9:1–6. <https://doi.org/10.1177/23794607241256709>.
66. Breiman L. Statistical Modeling: The Two Cultures (with comments and a rejoinder by the author). *Statistical Science* 2001;16. <https://doi.org/10.1214/ss/1009213726>.
67. Breiman L. Random forests. *Mach Learn* 2001;45:5–32. <https://doi.org/10.1023/A:1010933404324/METRICAL>.

68. Madhavan R, Kerr JA, Corcos AR, Isaacoff BP. Toward Trustworthy and Responsible Artificial Intelligence Policy Development. *IEEE Intell Syst* 2020;35:103–8. <https://doi.org/10.1109/MIS.2020.3019679>.
69. Duan W, Zhou S, Scalia MJ, Yin X, Weng N, Zhang R, et al. Understanding the Evolvement of Trust Over Time within Human-AI Teams. *Proc ACM Hum Comput Interact* 2024;8:1–31. <https://doi.org/10.1145/3687060>.
70. Fehlhaber AL, EL-Awad U. Trust development in online competitive game environments: a network analysis approach. *Appl Netw Sci* 2024;9:7. <https://doi.org/10.1007/s41109-024-00614-6>.

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.