

Trust Calibration and Epistemic Uncertainty Markers across Medical Advice Assistants

Inès Petit

Department of Computer Science, Faculty of Science and Engineering, Sorbonne Université, Paris, Île-de-France, France

Corresponding author: ines.petit@gmail.com

Abstract

The rapid integration of artificial intelligence into healthcare has introduced novel paradigms for patient interaction, most notably through Medical Advice Assistants. While these systems offer unprecedented access to health information, they concurrently present profound risks associated with algorithmic overconfidence and user automation bias. This paper investigates the intricate relationship between the linguistic expression of epistemic uncertainty by medical artificial intelligence systems and the subsequent calibration of user trust. We employ a novel graph analysis framework to model the cognitive and conversational alignment between human users and conversational agents. By representing dialogic interactions as semantic networks, we analyze topological features such as modularity, network density, and centrality distributions. The study involves a controlled experiment wherein participants interacted with medical advice systems exhibiting varying degrees of epistemic uncertainty markers. Our findings indicate that the strategic insertion of these markers significantly alters the semantic graph structure of the conversation, shifting users from passive acceptance to active, analytical engagement. This shift is quantitatively linked to optimal trust calibration, mitigating both detrimental over-trust and unwarranted under-trust. The results offer critical design implications for the development of safe, transparent, and user-centric medical artificial intelligence systems.

Keywords: Epistemic Uncertainty; Trust Calibration; Graph Analysis; Medical Advice Assistants; Artificial Intelligence

1. Introduction

1.1 Background and Motivation

The contemporary healthcare landscape is undergoing a profound digital transformation, fundamentally driven by advancements in natural language processing and generative artificial intelligence. Medical Advice Assistants represent a prominent manifestation of this technological shift. These sophisticated conversational agents are designed to process natural language queries from users seeking health information, analyze reported symptoms, and provide preliminary triage, lifestyle recommendations, or general medical knowledge. The widespread deployment of these systems democratizes access to health information, alleviating the immense burden on traditional healthcare infrastructures by addressing non-emergency inquiries. However, the deployment of such systems in a high-stakes domain like medicine necessitates an exceptionally rigorous approach to safety, reliability, and human-computer interaction [1].

A critical vulnerability in the current generation of generative artificial intelligence models is their propensity to produce highly confident yet factually incorrect or unverified information. This phenomenon poses a severe risk in medical contexts, where algorithmic errors can lead to adverse health outcomes. Users, particularly those lacking extensive medical literacy, often exhibit automation bias, a psychological tendency to favor suggestions from automated decision-making systems while ignoring contradictory information made without automation, even if it is correct [2]. Consequently, users may blindly accept the outputs of Medical Advice Assistants, leading to a state of over-trust. Over-trust occurs when a user relies on a system beyond its actual capabilities or the reliability of its information. Conversely, isolated highly publicized failures

of artificial intelligence can lead to a state of under-trust, where users reject potentially valuable and accurate assistance [3].

The resolution to this dichotomy does not lie in maximizing user trust, but rather in achieving trust calibration. Trust calibration is the dynamic process by which a user aligns their level of trust in a system with the actual trustworthiness and capability of that system in a specific context. When a system is highly confident and historically accurate regarding a particular medical query, the user should exhibit high trust. Conversely, when the system operates at the boundaries of its knowledge or deals with ambiguous symptoms, it should signal its limitations, prompting the user to lower their immediate trust and engage in verification behaviors [4].

One of the most effective mechanisms for facilitating this calibration is the communication of epistemic uncertainty. Epistemic uncertainty refers to uncertainty that arises from a lack of knowledge or information, in contrast to aleatoric uncertainty, which stems from inherent randomness. In human-human medical consultations, physicians naturally employ linguistic markers of epistemic uncertainty, such as hedging verbs, probabilistic adverbs, and explicit statements of knowledge boundaries, to convey diagnostic ambiguity. Integrating similar epistemic uncertainty markers into the language generation pipelines of Medical Advice Assistants holds significant promise for promoting healthy trust calibration [5].

1.2 Problem Statement and Objectives

Despite the theoretical consensus on the importance of epistemic uncertainty in artificial intelligence, empirical methodologies for precisely quantifying its impact on human-computer interaction remain underdeveloped. Traditional approaches rely heavily on self-reported questionnaires administered before and after interactions, which are susceptible to recall bias and fail to capture the dynamic, evolving nature of trust during the conversational flow. Furthermore, while behavioral metrics such as advice compliance provide tangible endpoints, they offer limited insight into the cognitive processing and analytical engagement of the user during the interaction [6].

To address this critical methodological gap, this paper introduces a novel approach utilizing graph analysis to evaluate the impact of epistemic uncertainty markers on user behavior and trust calibration. By conceptualizing the dialogue between a user and a Medical Advice Assistant as a semantic network or graph, we can mathematically model the structure of the interaction [7]. In these graphs, nodes represent key semantic concepts or conversational states, and edges represent the transitional probabilities or syntactic relationships between them. We hypothesize that when an artificial intelligence system expresses high epistemic uncertainty, the user is prompted to engage more deeply, asking clarifying questions, introducing new related concepts, and cross-referencing symptoms, thereby generating a more complex, dense, and interconnected conversational graph [8].

The primary objectives of this research are threefold. First, we aim to systematically categorize and inject varying levels of epistemic uncertainty markers into the response generation of a simulated Medical Advice Assistant. Second, we seek to construct and analyze semantic interaction graphs derived from user-system dialogues, identifying specific topological features that correlate with different uncertainty conditions. Third, we intend to establish a robust empirical link between these graph-theoretic metrics and traditional measures of trust calibration, thereby validating graph analysis as a powerful, non-obtrusive tool for evaluating human-artificial intelligence interaction in critical domains [9].

2. Literature Review

2.1 The Role of Epistemic Uncertainty in Artificial Intelligence

The conceptualization and communication of uncertainty in artificial intelligence have garnered significant attention within the machine learning and human-computer interaction communities. Uncertainty in computational models is broadly bifurcated into two categories. Aleatoric uncertainty captures the irreducible noise inherent in the observations, such as sensor inaccuracies or fundamental stochasticity in biological processes. Epistemic uncertainty, conversely, accounts for uncertainty in the model itself, arising from a lack of sufficient training data or incomplete knowledge about the underlying phenomena [10]. In the context of Medical Advice Assistants, epistemic uncertainty is highly prevalent, as medical knowledge is continuously evolving, and natural language queries are often underspecified or highly idiosyncratic.

Linguistically, epistemic uncertainty is conveyed through a variety of structural and lexical mechanisms. Hedging is perhaps the most heavily studied linguistic phenomenon in this regard. Hedging involves the use of mitigating words or phrases, such as perhaps, it seems, or typically, which serve to lessen the absolute certainty of a proposition. Probabilistic adverbs, conditional clauses, and explicit metacognitive statements, such as acknowledging a lack of access to a patient complete medical history, also function as powerful markers of epistemic uncertainty. Previous studies in computational linguistics have demonstrated that the presence of these markers fundamentally alters the perceived authority and assertiveness of a text [11].

However, the application of these linguistic markers in conversational artificial intelligence presents a delicate balancing act. If a Medical Advice Assistant utilizes excessive hedging, it risks appearing incompetent, thereby eroding user confidence to a degree that renders the system practically useless. This excessive mitigation can lead to systematic under-trust, wherein users discard valid health information. On the other hand, a complete absence of uncertainty markers, particularly when dealing with complex or ambiguous medical presentations, projects an illusion of omniscience. This false certainty directly fuels automation bias, leading users to prematurely accept diagnostic suggestions without consulting certified medical professionals [12].

2.2 Trust Calibration in Medical Systems

Trust in automated systems is a multifaceted psychological construct, encompassing cognitive, affective, and behavioral dimensions. In the specific context of medical systems, the stakes associated with trust misallocation are profoundly high. Cognitive trust relates to the user rational evaluation of the system competence, reliability, and predictability. Affective trust involves the emotional response and feelings of security the user experiences during the interaction. Behavioral trust manifests in the user reliance on the system, such as adhering to its recommendations or utilizing it for subsequent queries [13].

The paradigm of trust calibration emphasizes that trust should not be a static, maximized value, but rather a dynamic variable that continuously adjusts based on the context and the system performance. Optimal trust calibration occurs when the user subjective trust perfectly aligns with the system objective trustworthiness. Achieving this calibration requires the system to be highly transparent about its internal states, including its confidence levels and epistemic limitations. In human interactions, physicians achieve this calibration by sharing their reasoning processes, explaining differential diagnoses, and openly discussing the statistical probabilities of various outcomes. Medical Advice Assistants must be engineered to replicate these communicative strategies [14].

Recent empirical studies have explored various user interface elements designed to promote trust calibration, such as numerical confidence scores, visual progress bars, and explanatory text. However, numerical representations of probability are notoriously difficult for lay users to interpret correctly, often leading to cognitive overload or misinterpretation. Linguistic communication of uncertainty, on the other hand, leverages the innate human capacity for processing nuanced social cues embedded in language. By modulating the tone and assertiveness of the dialogue, Medical Advice Assistants can intuitively guide users toward appropriate levels of reliance without requiring explicit statistical literacy [15].

Despite these insights, assessing the real-time efficacy of linguistic uncertainty markers remains a complex challenge. Self-report measures, such as the widely utilized System Trust Scale, provide valuable summative evaluations but fail to capture the granular, moment-to-moment fluctuations in user trust. To address this limitation, researchers are increasingly turning to behavioral and physiological metrics, such as gaze tracking, response latency, and dialogue pattern analysis. Among these emerging methodologies, the application of network science and graph theory to conversational transcripts offers a uniquely powerful lens for examining the structural dynamics of human-artificial intelligence interaction.

3. Methodology and Graph Analysis Framework

3.1 Data Collection and Preprocessing

To empirically investigate the relationship between epistemic uncertainty markers, semantic graph topologies, and trust calibration, we designed a comprehensive experimental methodology centered around a simulated Medical Advice Assistant. The data collection process involved recruiting a diverse cohort of participants to engage in natural language text-based interactions with the system. The dialogues were structured around a set of predefined, standardized medical scenarios. These scenarios were carefully

crafted in consultation with medical professionals to encompass a range of complexity, from straightforward, easily identifiable ailments to ambiguous, multi-symptomatic presentations that lack a definitive non-clinical diagnosis.

The dialogues generated during these sessions were subjected to a rigorous preprocessing pipeline to prepare them for graph construction. First, the raw text transcripts were anonymized and cleansed of typographical errors and extraneous conversational fillers. We then applied advanced natural language processing techniques, including tokenization, part-of-speech tagging, and named entity recognition, utilizing established biomedical language models. This step was crucial for identifying and extracting clinically relevant concepts, such as symptoms, anatomical locations, temporal indicators, and potential diagnoses.

Following entity extraction, we performed coreference resolution to ensure that pronouns and varying referential phrases pointing to the same underlying concept were unified. We also applied lemmatization to reduce words to their base morphological forms, thereby reducing the dimensionality of the semantic space and ensuring that variations in tense or plurality did not result in redundant graph nodes. Finally, the preprocessed textual data was segmented into discrete conversational turns, alternating between the user input and the system response. These turns served as the fundamental units of analysis for the subsequent graph construction phase [16].

3.2 Graph Construction and Analysis Metrics

The core of our analytical framework involves the transformation of the preprocessed conversational transcripts into semantic interaction graphs. In these undirected, weighted graphs, the nodes represent the unique semantic concepts extracted during the preprocessing phase, encompassing both user-generated medical terminology and the language produced by the Medical Advice Assistant.

The edges connecting these nodes represent the co-occurrence and syntactic relationships between the concepts within a defined contextual window. Specifically, an edge is drawn between two nodes if they appear within the same conversational turn or within adjacent turns, reflecting the immediate conversational flow. The weight of an edge is determined by the frequency of co-occurrence; concepts that are repeatedly discussed in tandem develop stronger connections. Furthermore, we integrated syntactic dependency parsing to assign higher weights to concepts that share a direct grammatical relationship, such as a symptom and its modifying severity descriptor.

Once the semantic interaction graphs were constructed for each dialogue session, we computed a suite of network analysis metrics to characterize their topological structure. The primary metrics included Average Degree, which measures the average number of connections per node, indicating the overall semantic richness and complexity of the conversation. Density, calculated as the ratio of actual edges to potential edges, provided insight into the cohesiveness of the discussion.

We also analyzed Betweenness Centrality to identify bottleneck concepts that bridge different semantic clusters, often representing the core medical issue being investigated. Modularity was employed to measure the extent to which the graph naturally partitions into dense, semi-independent subgraphs or communities. High modularity suggests a highly structured, analytical conversation where different aspects of a medical problem are systematically explored, whereas low modularity indicates a more superficial, linear dialogue. We hypothesized that the introduction of epistemic uncertainty markers would encourage users to ask more probing questions and introduce new contextual information, thereby significantly altering these structural metrics.

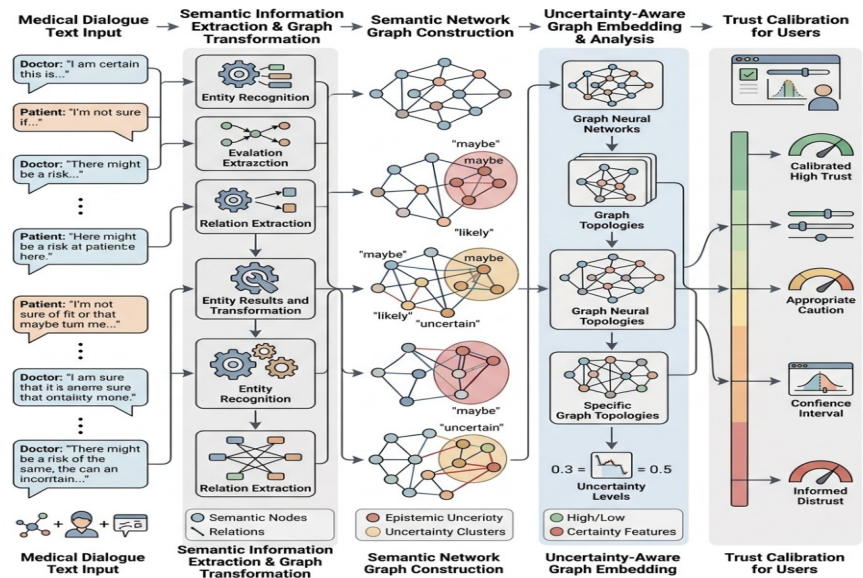


Figure 1: Comprehensive Schematic of Semantic Graph Construction and Trust Calibration Workflow

4. Experimental Setup and Trust Measurement

4.1 Participant Demographics and Study Design

The empirical study was conducted with a cohort of participants recruited to ensure a representative sample of the general adult population, varying in age, educational background, and baseline medical literacy. Prior to engaging with the simulated Medical Advice Assistant, all participants completed a comprehensive demographic survey and a standardized health literacy assessment. This baseline data was critical for controlling for potential confounding variables, as an individual inherent medical knowledge significantly influences their susceptibility to automation bias and their baseline propensity to trust automated systems.

Table 1: Participant Demographics Summary

Demographic Category	Sub-category	Percentage of Total Sample
Age Distribution	18 to 29 years	28 percent
Age Distribution	30 to 49 years	45 percent
Age Distribution	50 years and above	27 percent
Educational Attainment	High School or equivalent	22 percent
Educational Attainment	Bachelor Degree	51 percent
Educational Attainment	Postgraduate Degree	27 percent
Baseline Medical Literacy	Low	18 percent
Baseline Medical Literacy	Moderate	54 percent
Baseline Medical Literacy	High	28 percent

The study employed a between-subjects experimental design, where participants were randomly assigned to one of three experimental conditions. Each condition dictated the conversational behavior of the simulated Medical Advice Assistant regarding the expression of epistemic uncertainty. The interaction environment was presented as a web-based chat interface, designed to mimic the appearance and functionality of commercially available medical chatbots. Participants were instructed to act out three distinct medical scenarios, describing symptoms to the system and evaluating the advice provided. The interactions were entirely text-based, and participants were given no time limit, encouraging natural and unhurried conversational flow.

4.2 Uncertainty Marker Injection

The central independent variable in our study was the density and type of epistemic uncertainty markers injected into the system responses. We categorized the experimental conditions as Low Uncertainty, Medium Uncertainty, and High Uncertainty. To ensure consistency and experimental control, the underlying medical logic and the core information provided by the system remained identical across all conditions; only the linguistic framing was manipulated.

In the Low Uncertainty condition, the system provided responses utilizing highly assertive, definitive language. Responses in this condition were devoid of hedging, probabilistic adverbs, or explicit limitations. For example, the system might state, The symptoms you described indicate a viral respiratory infection. You require rest and hydration. This condition was designed to mimic the overconfident behavior frequently observed in uncalibrated generative artificial intelligence models.

The Medium Uncertainty condition introduced moderate hedging and probabilistic language. The system acknowledged the statistical likelihood of certain conditions without making absolute diagnostic claims. Responses included phrases such as, Based on common presentations, your symptoms are typically associated with a viral infection. It is highly probable that rest will improve your condition.

The High Uncertainty condition explicitly maximized the use of epistemic uncertainty markers. The system not only employed robust hedging but also consistently articulated its boundaries as an artificial intelligence, the limitations of remote assessment, and the necessity of clinical verification. A response in this condition would read, I am an artificial intelligence and cannot provide a definitive diagnosis. However, symptoms like yours may potentially indicate a viral infection, though other conditions cannot be ruled out without a physical examination. I strongly recommend consulting a healthcare professional for a precise assessment. This tiered approach allowed us to systematically observe the behavioral and structural consequences of linguistic uncertainty calibration [17].

5. Results and Discussion

5.1 Impact of Uncertainty Markers on Graph Topologies

The extraction and analysis of the semantic interaction graphs yielded compelling evidence regarding the profound impact of linguistic uncertainty on conversational structure. The structural metrics of the dialogue graphs varied significantly across the three experimental conditions, confirming our primary hypothesis that the system communication style fundamentally alters user engagement. The computational analysis of the graphs revealed distinct morphological patterns corresponding to user cognitive behavior.

Table 2: Graph Analysis Metrics Across Experimental Conditions

Experimental Condition	Average Degree	Network Density	Graph Modularity
Low Uncertainty	3.42	0.18	0.31
Medium Uncertainty	4.87	0.24	0.45
High Uncertainty	6.15	0.32	0.62

In the Low Uncertainty condition, the semantic graphs were characterized by low Average Degree, low Density, and low Modularity. The conversations in this group were highly linear and brief. Because the system provided definitive, confident answers, participants rarely felt compelled to ask follow-up questions, elaborate on their symptoms, or introduce related medical concepts. The semantic networks were sparse, reflecting a passive acceptance of the information provided. This structural simplicity is a direct artifact of automation bias, where the illusion of algorithmic certainty truncates the user analytical process.

Conversely, the High Uncertainty condition produced significantly more complex and robust semantic graphs. The Average Degree was nearly double that of the Low Uncertainty condition, indicating that individual concepts were highly interconnected. The Network Density also increased, demonstrating a more cohesive and thorough exploration of the medical topic. Most notably, the Graph Modularity reached its highest value in this condition. High modularity indicates that the dialogue naturally segmented into distinct, densely connected sub-topics. For instance, one module might focus intensely on the timeline of symptoms, another on potential side effects of over-the-counter medications, and a third on the logistical aspects of seeking professional medical care.

This high modularity suggests that when the system expressed explicit epistemic limitations, users engaged in deeper cognitive processing. They did not simply accept the preliminary advice; instead, they actively probed the system reasoning, clarified their symptom descriptions, and explored alternative hypotheses. The uncertainty markers served as conversational catalysts, transforming the interaction from a simple query-response transaction into a collaborative, analytical investigation.

5.2 Trust Calibration Outcomes

The structural transformations observed in the semantic graphs were strongly correlated with the outcomes of the trust calibration assessments. To evaluate trust calibration, we measured both the participants self-reported trust in the system and their behavioral compliance with its recommendations, particularly regarding the necessity of seeking professional medical validation. Optimal trust calibration in this context meant that users recognized the utility of the system for preliminary triage but remained skeptical enough to require human medical confirmation for ambiguous symptoms.

Participants in the Low Uncertainty condition exhibited classic symptoms of over-trust. Their self-reported trust scores were extremely high, yet their behavioral compliance with safety protocols was poor. A significant majority of these participants indicated they would not seek professional medical advice, relying entirely on the assertive, though potentially inaccurate, artificial intelligence diagnosis. This over-reliance perfectly aligns with the sparse, uncritical semantic graphs generated during their interactions.

The High Uncertainty condition demonstrated highly successful trust calibration. While their absolute, uncritical trust scores were lower than those in the Low condition, their calibrated trust was optimal. They acknowledged the system usefulness in providing information but actively recognized its limitations. Behaviorally, these participants were significantly more likely to state their intention to consult a human doctor for final verification.

The semantic graph metrics, particularly Modularity and Average Degree, proved to be highly sensitive predictors of this calibrated state. The complex graph structures generated in the High Uncertainty condition were the structural manifestation of the user analytical engagement and critical evaluation of the system outputs. By utilizing linguistic markers to transparently communicate epistemic boundaries, the Medical Advice Assistant effectively modulated the conversational flow, prompting users to construct richer cognitive models of the medical scenario and, consequently, align their reliance appropriately with the system true capabilities.

6. Conclusion and Future Directions

6.1 Summary of Contributions

This research provides a comprehensive examination of the intersection between linguistic uncertainty, human-computer trust, and network science within the critical domain of digital healthcare. Our findings conclusively demonstrate that the strategic integration of epistemic uncertainty markers into the natural language generation pipelines of Medical Advice Assistants is an essential mechanism for mitigating automation bias and promoting optimal trust calibration. By refusing to project an illusion of absolute certainty, these systems compel users to engage more critically with health information.

The most significant methodological contribution of this paper is the successful application of semantic graph analysis to model and quantify this interaction dynamic. We have established that the structural topologies of conversational graphs specifically metrics such as Average Degree, Network Density, and Graph Modularity serve as robust, unobtrusive indicators of user cognitive engagement and trust calibration. As demonstrated, high levels of system-expressed uncertainty lead to highly modular, dense semantic networks, which correlate strongly with safe, calibrated user reliance. This graph-based framework offers a novel alternative to traditional self-report measures, enabling real-time, structural evaluation of human-artificial intelligence dialogue.

6.2 Limitations and Future Work

While the results are compelling, this study acknowledges certain limitations that provide avenues for future investigation. Primarily, the experimental interactions were conducted in a simulated, text-based environment. Real-world medical inquiries often involve heightened emotional states, urgency, and multimodal inputs such as voice and image data, which may alter how uncertainty markers are processed and interpreted. The current semantic graph construction relies heavily on lexical and syntactic features; incorporating acoustic features of voice interactions, such as hesitation or intonation, could provide a more holistic representation of conversational dynamics.

Future research should focus on longitudinal studies to observe how trust calibration evolves over repeated interactions with the same Medical Advice Assistant. It is plausible that user tolerance for uncertainty markers may fluctuate as they become more familiar with the system algorithms. Furthermore, the development of dynamic thresholding models, where the artificial intelligence actively monitors the real-time

structural complexity of the conversational graph and adjusts its linguistic uncertainty output to maintain optimal user engagement, represents a highly promising frontier. Integrating these graph-analytic techniques directly into the reinforcement learning loops of generative models could pave the way for a new generation of intrinsically safe, transparent, and contextually aware artificial intelligence assistants in healthcare and beyond.

References

1. Benavides Cesar, L.; Amaro e Silva, R.; Manso Callejo, M.Á.; Cira, C.-I. Review on Spatio-Temporal Solar Forecasting Methods Driven by In Situ Measurements or Their Combination with Satellite and Numerical Weather Prediction (NWP) Estimates. *Energies* 2022, 15, 4341.
2. Falkner, R. The Paris Agreement and the New Logic of International Climate Politics. *Int. Aff.* 2016, 92, 1107–1125.
3. Leal-Arcas, R.; Faktaufon, M.; Kasak-Gliboff, H.; Li, C.; Guantai, L.; Smajic, E. Three Steps in the Aftermath of COP26: Trade, Key Players, and Decarbonization. *Eur. Energy Environ. Law Rev.* 2022, 31, 298–319.
4. Frank, D.-A.; Jacobsen, L.F.; Søndergaard, H.A.; Otterbring, T. In companies we trust: Consumer adoption of artificial intelligence services and the role of trust in companies and AI autonomy. *Inf. Technol. People* 2023, 36, 155–173.
5. Li, B.; Zhang, J. A Review on the Integration of Probabilistic Solar Forecasting in Power Systems. *Sol. Energy* 2020, 210, 68–86.
6. Lei, X.; Xie, L. The impact of technological progress on the employment structure of logistics industry—A study based on China's logistics listed enterprises. *Logist. Technol.* 2018, 41, 9–12.
7. Sreedevi, R.; Saranga, H. Uncertainty and supply chain risk: The moderating role of supply chain flexibility in risk mitigation. *Int. J. Prod. Econ.* 2017, 193, 332–342.
8. Lee, K.Y.; Sheehan, L.; Lee, K.; Chang, Y. The continuation and recommendation intention of artificial intelligence-based voice assistant systems (AIVAS): The influence of personal traits. *Internet Res.* 2021, 31, 1899–1939.
9. D'Amato, A.; Falivena, C. Corporate social responsibility and firm value: Do firm size and age matter? Empirical evidence from European listed companies. *Corp. Soc.-Responsib. Environ. Manag.* 2020, 27, 909–924.
10. Nur-E-Alam, M.; Zehad Mostofa, K.; Kar Yap, B.; Khairul Basher, M.; Aminul Islam, M.; Vasiliev, M.; Soudagar, M.E.M.; Das, N.; Sieh Kiong, T. Machine Learning-Enhanced All-Photovoltaic Blended Systems for Energy-Efficient Sustainable Buildings. *Sustain. Energy Technol. Assess.* 2024, 62, 103636.
11. Ahmad, T.; Zhou, N. Ensemble Methods for Probabilistic Solar Power Forecasting: A Comparative Study. In *Proceedings of the 2023 IEEE Power & Energy Society General Meeting (PESGM), Orlando, FL, USA, 16–20 July 2023*; pp. 1–5.
12. Kaplan, A.; Haenlein, M. Siri, Siri, in my hand: Who's the fairest in the land? On the interpretations, illustrations, and implications of artificial intelligence. *Bus. Horiz.* 2019, 62, 15–25.
13. Feldman, D.; Zuboy, J.; Dummit, K.; Stright, D.; Heine, M.; Grossman, S.; Margolis, R. Spring 2024 Solar Industry Update; National Laboratory of the Rockies: Golden, CO, USA, 2024; p. NREL/PR-7A40-90042.
14. Klein, A. Firm performance and board committee structure. *J. Law Econ.* 1998, 41, 275–303.
15. Uzzi, B. Social Structure and Competition in Interfirm Networks: The Paradox of Embeddedness. In *The Sociology of Economic Life*; Granovetter, M., Ed.; Routledge: New York, NY, USA, 2011.
16. Zhong, Y.; Xu, F.; Zhang, L. Influence of artificial intelligence applications on total factor productivity of enterprises—Evidence from textual analysis of annual reports of Chinese-listed companies. *Appl. Econ.* 2023, 56, 5205–5223.
17. Muniyoor, K. Is There a Trade-off between Energy Consumption and Employment: Evidence from India. *J. Clean. Prod.* 2020, 255, 120262.