



Grund Institute

13th January 2026

Evaluating Human–AI Interaction Styles: Dimensions, Measurement, and Framework

The ways individuals respond to probabilistic AI guidance

Prepared by:

Dr. Sanela Vellino, Ph.D

Grund Institute, Switzerland



Introduction

Recent advances in AI (e.g. conversational agents) make it crucial to understand how individuals respond to probabilistic AI guidance. Research shows that familiarity and past use of AI tends to increase trust and acceptance of new AI applications[1]. However, underlying cognitive and personality factors (such as ambiguity tolerance or need for control) also shape these responses. For example, people who **tolerate ambiguity** (discomfort with uncertainty) tend to feel less fear and more acceptance of AI systems[2]. Individuals differ in **control preference** (whether they want the system to take initiative or prefer to steer the process). In one study of conversational AI, a “control preference” trait (measured on a Likert scale) predicted whether users preferred system- vs. user-initiated interactions[3]. Similarly, baseline **trust in algorithms** varies by person: many studies find people generally trust human judgment more than algorithmic decisions, exhibiting “algorithm aversion” unless trust is calibrated[4][5]. Finally, **cognitive engagement style** (e.g. *need for cognition* or analytic thinking) influences how users process AI output. Those who enjoy effortful reasoning (high need for cognition) are more likely to scrutinize or effectively use AI advice[6][7].

- **Tolerance for Ambiguity:** This cognitive trait reflects comfort with uncertain or incomplete information. In cognitive science, ambiguity tolerance relates to how much one prefers structured, clear data vs. exploring uncertain scenarios. Empirical evidence shows that individuals with low ambiguity tolerance tend to view AI’s probabilistic outputs negatively, whereas those with higher tolerance (and a preference for complexity/novelty) report more positive attitudes toward AI[2]. This aligns with findings that discomfort with ambiguity predicts AI *fear*, while a need for complexity/novelty predicts AI *acceptance*[2].

- **Control Preference:** Stemming from behavioral psychology and organizational theory (e.g. locus-of-control, autonomy needs), this dimension measures whether a person prefers to direct actions or let the system propose options. For instance, in dialogue-based AI, users vary in wanting the system to initiate topics vs. those who only respond when prompted. One recent study operationalized this on a 5-point scale and found it significantly shaped user choice of interactive style[3]. People with high control preference may require more interactive control knobs or approvals when using AI.
- **Trust Calibration:** This social-psychological dimension gauges whether a user's trust in AI is appropriately aligned with system reliability. Classic trust research defines trust as “the attitude that an agent will help achieve goals in uncertainty”[8]. Behavioral studies show that most people under-trust algorithmic advice by default: for example, participants trusted a human recruiter's decision more than an equivalent algorithm's decision[4]. Algorithmic decision-making is often judged as less fair or trustworthy unless transparency and accuracy are highlighted[4][5]. Thus “trust calibration” (avoiding both undue skepticism and overtrust) is a key factor in AI interaction style.
- **Cognitive Engagement (Need for Cognition / Reflection):** Borrowed from cognitive psychology, this trait measures how much a person enjoys and engages in analytic thinking. High-need-for-cognition individuals relish complex problem-solving and are likelier to consider multiple hypotheses. Research in human–AI collaboration (e.g. among physicians using AI support) has identified *need for cognition* as influencing how rigorously people vet AI advice[6]. Likewise, studies of algorithm trust cite *cognitive reflection* (analytic vs. intuitive decision style) as a moderator of transparency's effect on trust[7]. These findings imply that cognitive style (analytic vs heuristic) affects how one interprets probabilistic outputs and explanatory information from AI.

Psychometric Foundations of the Four Interaction Types

The four human–AI interaction types presented in this paper are grounded in established psychometric and behavioral science principles. Rather than relying on self-reported attitudes or surface-level usage patterns, the framework is derived from a structured measurement approach designed to capture **latent cognitive interaction tendencies**.

At a high level, the development of the diagnostic follows recognized best practices in psychological and behavioral measurement. These practices are widely used in research settings to ensure that constructs are meaningful, stable, and empirically distinguishable, while remaining appropriate for real-world application.

The development process includes:

1. **Construct Definition and Conceptual Modeling**
Core cognitive dimensions relevant to interaction with probabilistic systems are defined based on existing research in cognitive psychology, decision science, and human–AI interaction. These dimensions are conceptual rather than behavioral checklists, ensuring that the framework captures underlying interaction tendencies rather than momentary preferences.
2. **Iterative Instrument Development**
Measurement prompts are developed and refined through expert review and qualitative

testing to ensure they reflect the intended cognitive constructs without signaling preferred responses. This step focuses on conceptual clarity rather than item disclosure.

3. **Empirical Structure Testing**

Pilot data is used to examine whether observed response patterns align with the intended conceptual structure, and whether dimensions remain distinct from one another when tested empirically. This ensures that the framework reflects real cognitive variation rather than theoretical assumptions.

4. **Reliability and Stability Assessment**

The framework is evaluated for internal consistency and stability over time, ensuring that identified interaction tendencies are not driven by situational noise or transient states.

5. **Measurement Calibration**

Advanced measurement techniques are applied to ensure that results are comparable across individuals and contexts, and that the framework captures meaningful differences along each cognitive dimension rather than binary classifications.

6. **Validity Evaluation**

The framework is assessed against external indicators, including observed behavior in AI-supported tasks, to confirm that identified interaction types correspond to real differences in how individuals engage with probabilistic systems.

7. **Continuous Refinement**

As the framework is applied across different organizational contexts, it is iteratively refined to improve robustness, boundary definition, and applicability, without altering the core conceptual structure.

This approach mirrors established standards in psychometric scale development, including construct definition, iterative testing, empirical validation, and continuous refinement [9][10][11]. Importantly, while the scientific foundations of the framework are transparent, the specific implementation details, scoring logic, and classification thresholds remain intentionally non-public to preserve methodological integrity and prevent misuse.

The result is a diagnostic framework that enables reliable classification of human–AI interaction types without reducing individuals to traits, scores, or rankings, and without exposing proprietary measurement logic.

Comparison with Traditional Instruments

Unlike personality inventories or strengths assessments (which measure broad, context-free traits), this AI-interaction tool targets specific, context-dependent preferences. Traditional tools like the Big Five or behavioral profiles might label someone as “analytical” or “extroverted,” but they do not directly tell us how that person handles probabilistic AI output. For example, two people high in *openness* might still differ sharply in their trust of algorithms – a nuance captured only by an AI-specific trust scale, not by general personality[4]. Strengths-based diagnostics identify core talents (e.g. creativity, decisiveness) but do not assess technology-specific attitudes. Behavioral style tests (e.g. conflict or learning styles) similarly focus on interpersonal or learning preferences rather than how one processes ambiguous data.

In contrast, our proposed scales are designed for the **human–AI context**: they measure how a person **actually interacts with an algorithm**, not just how they see themselves generally. This makes the tool more predictive of AI-relevant outcomes (like reliance on AI vs. double-checking it) than general assessments. Indeed, research shows that algorithmic trust and reliance depend on task specifics and user state (e.g. explained in cultural dimension models[5][4]) – factors outside standard personality tests. Thus, while existing instruments have value for general development, they offer limited insight into AI collaboration without an AI-centric framework.

Human–AI Collaboration: Cognitive Load and Trust

Cognitive Load in Probabilistic Reasoning: Cognitive science shows that probabilistic reasoning is mentally taxing. Studies on Bayesian reasoning tasks reveal that individuals with higher working memory capacity make significantly fewer reasoning errors[12]. Conversely, people with lower working memory struggle with abstract probabilities unless supported. For instance, presenting statistical information as icon arrays or visual aids dramatically improves accuracy for low-capacity individuals, reducing cognitive load[13]. This implies that AI systems presenting probabilistic forecasts should tailor formats (text vs. graphics) to the user’s cognitive style. It also suggests a dimension of interaction style: some users need highly interpretable representations, while others can directly handle numeric outputs.

Trust and Algorithm Aversion: Behavioral research consistently finds that even highly accurate algorithms may be mistrusted. A systematic review on algorithm aversion notes that people often resist algorithmic decisions “consciously or unconsciously” despite better performance[5]. In personnel selection experiments, for example, algorithmic and human choices were equated, yet participants rated the human-taken action as more trustworthy and acceptable[4]. This highlights how human teams will see AI advice through a skeptical lens unless trust is actively managed. Models like the CHAI-T framework emphasize that trust in AI is dynamic and multi-faceted[14]: it depends on user traits, system transparency, and evolving collaboration context. For instance, giving clear explanations can improve trust for some users (cognitive reflection interacts here[7]), whereas others might need consistent accuracy signals.

Adaptive Collaboration: Cognitive science and HCI suggest that optimal human–AI collaboration requires alignment (cognitive fit) between the system’s behavior and the user’s needs. Zhao et al. propose that AI should *learn* each user’s decision patterns and adapt its communication style accordingly[15]. An AI that adjusts its level of autonomy, explanation detail, or framing of uncertainty based on a user’s profile can greatly improve teamwork outcomes. For example, an “analytical” user might benefit when the AI provides raw probability distributions, whereas a “cautious” user might need deterministic thresholds or more interactive control.

Cognitive Offloading: Finally, recent evidence warns that over-reliance on AI may erode human skills. In an educational context, heavy AI tool use was linked to lower critical thinking scores (partly via cognitive offloading)[16]. This suggests that interaction styles also have learning consequences: users who defer too much to AI might atrophy their own reasoning, while those who engage critically maintain sharper skills. A diagnostic that flags “overly dependent” vs. “appropriately independent” styles could thus inform training to prevent cognitive atrophy.

Four Human–AI Interaction Archetypes

Combining these dimensions suggests four prototypical user *styles*, akin to quadrant models in decision theory[17]. For illustration, consider a two-axis grid of (1) **Trust in AI** (high vs. low) and (2) **Ambiguity Tolerance** (high vs. low). This yields four archetypes:

- **Confident Explorer** (High Trust, High Tolerance): Embraces AI guidance and is comfortable with its inherent uncertainty. These users readily delegate tasks to AI, are open to novel suggestions, and often see AI as a creative partner. They tend to engage deeply (high need for cognition) and will apply AI insights broadly, knowing they can handle imperfect results. (Optimal collaboration occurs when the AI adapts creativity and depth of information to match their exploratory mindset.)

- **Prudent Analyst** (Low Trust, High Tolerance): Comfortable with data ambiguity but initially skeptical of AI. They will rigorously analyze any AI output, double-checking results against logic or data. Their high cognitive engagement means they can extract value from probabilistic suggestions, but only if convinced of reliability. In practice, they benefit from transparent AI (e.g. clear confidence intervals or evidence citations) to calibrate trust before use.
- **Structured Controller** (High Trust, Low Tolerance): Generally willing to rely on AI but needs clear, deterministic guidance. They prefer outputs that minimize ambiguity (e.g. binary recommendations or precise probabilities with “high-confidence” labels). These users like to maintain oversight via checkpoints or thresholds. They trust the AI’s competence but require step-by-step or rule-based interaction, reflecting a strong need for control structure.
- **Cautious Traditionalist** (Low Trust, Low Tolerance): Uncomfortable with uncertainty and unwilling to blindly trust AI. They use AI only sparingly, typically under high-certainty conditions. More likely to stick with familiar (often human-provided) information, they view probabilistic outputs as noise. To engage them, an AI must prove its worth by behaving deterministically (e.g. frequent accuracy feedback, robust validation).

These categories mirror classic decision-style taxonomies that use similar axes (e.g. directive vs. analytic) to yield four styles[17]. Importantly, the styles are not fixed “types” but idealized clusters: most people show a dominant style with some flexibility. Framing users this way helps target training and AI interface design. For instance, a Confident Explorer might be encouraged to “push boundaries” with advanced AI features, while a Structured Controller might receive a more guided, transparent interface.

Conclusion

In summary, a comprehensive **human–AI style framework** combines cognitive science (ambiguity tolerance, cognitive load), psychology (trust, need for control), and organizational theory (adaptation to roles) to form the basis of four interaction archetypes. A psychometrically sound diagnostic would measure these underlying traits via multi-item scales calibrated with factor analysis and IRT[11][10]. Unlike generic personality tests, this tool is specialized for AI collaboration contexts[4][17]. Empirical studies of human-AI teams underline the need for such profiling: by identifying where individuals fall on trust and tolerance spectra, we can match AI interfaces and training to their cognitive fit. In practice, this means managers and designers can predict which users need more explanation, which will intuitively leverage AI, and so on, ultimately unlocking better AI augmentation across the organization.

Sources: Concepts and evidence are drawn from cognitive psychology of uncertainty[2][12], behavioral studies of algorithm aversion and trust[5][4], human–AI collaboration research[15][14], and psychometric methodology texts[11][10], among others.

References:

[1] Adopting AI: how familiarity breeds both trust and contempt - PMC

<https://pmc.ncbi.nlm.nih.gov/articles/PMC10175926/>

[2] Associating attitudes towards AI and ambiguity: The distinction of acceptance and fear of AI - ScienceDirect

<https://www.sciencedirect.com/science/article/pii/S0001691825008947>

[3] Understanding User Preferences for Interaction Styles in Conversational Recommender Systems: The Predictive Role of System Qualities, User Experience, and Traits

<https://www.arxiv.org/pdf/2508.02328>

[4] (PDF) People's Reactions to Decisions by Human vs. Algorithmic Decision-makers: The Role of Explanations and Type of Selection Tests

https://www.researchgate.net/publication/364196290_People's_Reactions_to_Decisions_by_Human_vs_Algorithmic_Decision-makers_The_Role_of_Explanations_and_Type_of_Selection_Tests

[5] What influences algorithmic decision-making? A systematic literature review on algorithm aversion - ScienceDirect

<https://www.sciencedirect.com/science/article/pii/S0040162521008210>

[6] [8] Psychological Factors Influencing Appropriate Reliance on AI-enabled Clinical Decision Support Systems: Experimental Web-Based Study Among Dermatologists - PMC

<https://pmc.ncbi.nlm.nih.gov/articles/PMC12008695/>

[7] Microsoft Word - CHERNOSKUTOVA_2062656_MA-THESIS.docx

<http://arno.uvt.nl/show.cgi?fid=154703>

[9] Item response theory for measurement validity - PMC

<https://pmc.ncbi.nlm.nih.gov/articles/PMC4118016/>

[10] Artificial Intelligence Attitudes Inventory (AIAI): development and validation using Rasch methodology | Current Psychology

<https://link.springer.com/article/10.1007/s12144-025-08009-1>

[11] Frontiers | Best Practices for Developing and Validating Scales for Health, Social, and Behavioral Research: A Primer

<https://www.frontiersin.org/journals/public-health/articles/10.3389/fpubh.2018.00149/full>

[12] Frontiers | The Effects of Working Memory and Probability Format on Bayesian Reasoning

<https://www.frontiersin.org/journals/psychology/articles/10.3389/fpsyg.2020.00863/full>

[13] [2302.00707] Why Combining Text and Visualization Could Improve Bayesian Reasoning: A Cognitive Load Perspective

<https://arxiv.org/abs/2302.00707>

[14] Collaborative human-AI trust (CHAI-T): A process framework for active management of trust in human-AI collaboration - ScienceDirect

<https://www.sciencedirect.com/science/article/pii/S2949882125000842>

[15] The Role of Adaptation in Collective Human–AI Teaming - PMC

<https://pmc.ncbi.nlm.nih.gov/articles/PMC12093936/>

[16] AI Tools in Society: Impacts on Cognitive Offloading and the Future of Critical Thinking

<https://www.mdpi.com/2075-4698/15/1/6>

[17] regent.edu

<https://www.regent.edu/wp-content/uploads/2023/05/Regent-Readiness-Inventory-Decision-Making-Decision-Style-Inventory.pdf>