

When AI Changes its Tone, Does Acceptance Follow?

Kevin Lejeune¹ and Corentin Burnay²

¹Department of Strategic Support, CHU UCL Namur, Belgium and NADI - MINDIT Research Center, University of Namur, Belgium

²Department of Business administration and Management, NADI - MINDIT Research Center, University of Namur, Belgium

kevin.lejeune@unamur.be

corentin.burnay@unamur.be

Abstract: The integration of artificial intelligence (AI) into decision-making processes raises numerous questions about acceptance, particularly in hospital environments marked by strong professional identities. This working paper presents an ongoing experimental study that explores how the conversational tone of an AI agent might influence the acceptance of its recommendations, depending on the psychological profile of the healthcare professional. Building on a 2x2x2 model structured around three dimensions of ego (personal value, perceived competence and social role), the study aims to demonstrate that adapting the tone can significantly improve acceptance to algorithmic recommendations. The experimental protocol involves profession-specific scenarios in which participants assess a series of AI-generated messages, each introduced during the option evaluation stage of the decision-making process. By proposing a novel identity-based approach to AI design, this study seeks to contribute to both theory and practice. It is expected to open perspectives for the development of conversational agents capable of dynamically adjusting their tone to enhance acceptability and support professional autonomy in hospital settings.

Keywords: Conversational AI, Ego profiles, Decision-Making process, AI recommendation acceptance, Hospital professionals

1. Introduction

Artificial intelligence (AI) is gradually transforming decision-making practices in healthcare organizations, providing algorithmic support that is both powerful and fast (An et al, 2023; Davenport & Kalakota, 2019). However, this technological promise faces a more nuanced reality: the acceptance of AI-generated recommendations by healthcare professionals remains partial, uneven, and often silently resistant (Coeckelbergh, 2020; Martineau & Godin, 2023).

A substantial body of literature has examined the conditions for technology adoption, notably through models such as TAM (Davis, 1989), UTAUT (Venkatesh et al, 2003) and UTAUT2 (Venkatesh et al, 2012). These approaches explain adoption in terms of perceived usefulness, ease of use, or trust in the system. Yet, they do not always account for the identity-related specificities of hospital professionals, nor the interactional dimension of certain types of AI, such as conversational recommendation agents.

Our research aims to address this gap by investigating a central hypothesis: the acceptance of an AI recommendation depends not only on its content or source, but also on its tone; and how that tone resonates with the recipient's psychological profile. This represents an original theoretical and empirical contribution, seeking to better understand the mechanisms of interaction between humans and intelligent systems in contexts where professional identity plays a central role in stability, legitimacy, and decision-making power.

2. Conceptual Framework

2.1 Conversational Recommendation AI in Hospital Settings

This study focuses on a specific form of AI: the conversational recommendation agent. Such systems deliver suggestions in natural language to support professional reasoning. Unlike rule-based engines or silent interfaces, they initiate exchanges sometimes perceived as “human-like”, distinguishing them from other interaction modes. Rather than being prescriptive or autonomous, they belong to the category of “augmented” or “collaborative AI” (Glikson & Woolley, 2020), designed to enhance, rather than replace, human analytical capabilities.

In our protocol, the AI intervenes during a key stage of the decision-making process: option evaluation, situated between problem identification and final decision (Simon, 1960). This moment of subjective arbitration involves complex interactions between professional judgment, self-image, and message formulation. The importance of this stage stems from the fact that the healthcare professional is not yet

committed to action but is situated in a cognitive deliberation space, where the tone of a message can subtly, yet decisively, shape how its content is received.

2.2 Ego as a Moderating Variable

Professional ego is conceived here as a relatively stable set of internal representations, structured around three components: *personal value*, defined as the recognition of one's own professional legitimacy; *perceived competence*, referring to the confidence an individual has in their ability to act effectively; and *social role*, understood as the perceived position within a hierarchy or institution. Unlike *identity threat* (Selenko et al, 2022), a contextual reaction triggered by the introduction of a new technology, ego here is a pre-existing structural filter, present before any interaction with the AI.

To explain reactions to tone-ego mismatches, three psychological mechanisms are mobilized. *Cognitive dissonance* (Festinger, 1957) helps explain how a message contradicting one's positive self-image may lead to rejection. *Threat to autonomy* (Deci & Ryan, 2000) arises when a message is seen as overly directive, potentially restricting freedom of action, especially for professionals with high perceived competence. Lastly, *system justification* (Jost, 2020; Jost & Banaji, 1994) sheds light on defensive or identity-based reactions when a de-hierarchized message is interpreted as delegitimizing institutional status. These established mechanisms in the social sciences offer a novel framework for linking tone preferences to identity dynamics, especially salient in hospital environments.

2.3 The 2x2x2 Model: Ego Profiles and Tone Hypotheses

Combining two levels (low/high) for each of the three ego dimensions yields eight theoretical psychological profiles. Each is associated with a theoretical preference for AI message tone. For instance, a professional with high personal value, high perceived competence, and a strong social role might respond positively to a collaborative and affirming tone, but negatively to a neutral, impersonal tone.

The central hypothesis is that alignment between message tone and ego profile maximizes recommendation acceptance. Eight hypotheses (H1 to H8) have been formulated, each corresponding to a specific profile and its preferred tone. This framework enables the prediction of optimal interactional conditions to foster the acceptance of AI-generated messages.

3. Methodology

3.1 Experimental Design and Scenario Preparation

The study follows a two-phase protocol. The first phase establishes each participant's psychological profile using a structured questionnaire composed of nine items, equally covering the three ego dimensions. Personal value is assessed via an adaptation of the Rosenberg Self-Esteem Scale (Rosenberg, 1965), focused on professional self-esteem; perceived competence through the General Self-Efficacy Scale (Schwarzer & Jerusalem, 1995); and social role using an adapted Subjective Social Status Scale (Adler et al, 2000), measuring perceived recognition within the organization.

In the second phase, participants review AI-generated messages identical in content but differing in tone. Eight messages correspond to profiles from the 2x2x2 model, while a ninth, neutral-toned message serves as a baseline. These messages are presented within profession-specific scenarios, selected from a set of six realistic options prepared for each professional group by a panel of experts. The AI message is introduced during the option evaluation stage of the decision-making process. It is not intended to prescribe a course of action, but rather to support professional reasoning. This ensures both contextual realism and a clear distinction from systems that operate autonomously or impose decisions. Participants rate each message by answering: "To what extent does this message seem appropriate and worthy of being considered in your professional reasoning?", using a seven-point Likert scale (1 = not at all, 7 = completely).

3.2 Planned Analysis

Data analysis will begin with an ANOVA to test interactions between ego profiles and message tones. This initial step aims to empirically verify whether certain combinations are statistically associated with higher acceptance. A second step will involve regression models to estimate the specific contribution of each ego dimension to the level of acceptance. This will help determine whether certain dimensions have a more structuring effect on message reception than others. Finally, classification analyses will explore the influence of control variables (profession, experience, gender, age, etc.) on the expressed preferences. The objective is to assess whether contextual or professional status effects modulate the observed results.

This protocol isolates tone effects while accounting for interindividual variability. It also opens the door to identifying atypical profiles or counterintuitive configurations that could enrich theoretical modeling and inspire new directions in the conversational design of AI systems in hospital environments.

4. Expected Discussion

This study could lead to a psychological segmentation of AI users in hospital settings. It would suggest that, beyond technical aspects (such as usefulness or explainability), message reception may depend on tone and its alignment with the recipient's identity. Such findings would complement work on system personalization (Seeber et al., 2020) and highlight the potential of ego-based adaptive design. It might also raise important questions about the role of language as a symbolic mediator in human-AI relationships, an aspect still largely underexplored in research dominated by functionalist perspectives.

5. Future Perspectives

Beyond the empirical validation of the proposed model, this research could open avenues for developing AI systems capable of adjusting their tone to user's psychological profiles, in the interest of ethical acceptability and effective support. It would also invite reflection on which stages of the decision-making process AI should be involved in, without undermining professional autonomy. Ultimately, these results might inform the design of health systems that are more attuned to identity dynamics, enhancing their effectiveness while respecting the human and institutional balance specific to hospital environments.

Ethics Declaration: Ethical approval was not required for the research presented in this paper.

AI Declaration: AI was employed to generate the simulated recommendation messages used in the experimental phase, based on predefined tone variations. As the study focuses on the impact of such messages, its use was intentional and aligned with the research objectives. All conceptual design and theoretical framework were entirely developed by the authors.

References

- Adler, N.E., Epel, E.S., Castellazzo, G. and Ickovics, J.R. (2000) "Relationship of subjective and objective social status with psychological and physiological functioning: Preliminary data in healthy, White women", *Health Psychology*, Vol 19, No. 6, pp. 586–592. <https://doi.org/10.1037/0278-6133.19.6.586>
- An, A., Rahman, M.S., Zhou, J. and Kang, J.J. (2023) "A comprehensive review on machine learning in healthcare industry: Classification, restrictions, opportunities and challenges", *Sensors*, Vol 23, No. 9, 4178. <https://doi.org/10.3390/s23094178>
- Coeckelbergh, M. (2020) *AI Ethics*, MIT Press.
- Davenport, T.H. and Kalakota, R. (2019) "The potential for artificial intelligence in healthcare", *Future Healthcare Journal*, Vol 6, No. 2, pp. 94–98. <https://doi.org/10.7861/futurehosp.6-2-94>
- Davis, F.D. (1989) "Perceived usefulness, perceived ease of use, and user acceptance of information technology", *MIS Quarterly*, Vol 13, No. 3, pp. 319–340. <https://doi.org/10.2307/249008>
- Deci, E.L. and Ryan, R.M. (2000) "The "what" and "why" of goal pursuits: Human needs and the self-determination of behavior", *Psychological Inquiry*, Vol 11, No. 4, pp. 227–268. https://doi.org/10.1207/S15327965PLI1104_01
- Festinger, L. (1957) *A theory of cognitive dissonance*, Stanford University Press.
- Glikson, E. and Woolley, A.W. (2020) "Human trust in artificial intelligence: Review of empirical research", *Academy of Management Annals*, Vol 14, No. 2, pp. 627–660. <https://doi.org/10.5465/annals.2018.0057>
- Jost, J.T. (2020) *A Theory of System Justification*, Harvard University Press.
- Jost, J.T. and Banaji, M.R. (1994) "The role of stereotyping in system-justification and the production of false consciousness", *British Journal of Social Psychology*, Vol 33, No. 1, pp. 1–27. <https://doi.org/10.1111/j.2044-8309.1994.tb01008.x>
- Martineau, J.T. and Godin, F.R. (2023) "Tour d'horizon des enjeux éthiques liés à l'IA en santé", *Éthique publique*, Vol 25, No. 1. <https://doi.org/10.4000/ethiquepublique.7978>
- Rosenberg, M. (1965) *Society and the adolescent self-image*, Princeton University Press. <https://doi.org/10.1515/9781400876136>
- Schwarzer, R. and Jerusalem, M. (1995) "Generalized Self-Efficacy Scale", in Weinman, J., Wright, S. and Johnston, M. (eds.) *Measures in health psychology: A user's portfolio. Causal and control beliefs*. Windsor, UK: NFER-NELSON, pp. 35–37.
- Seeber, I., Bittner, E.A.C., Briggs, R.O., de Vreede, T., de Vreede, G.-J., Elkins, A., Maier, R. and Merz, A.B. (2020) "Machines as teammates: A research agenda on AI in team collaboration", *Information & Management*, Vol 57, No. 2, 103174. <https://doi.org/10.1016/j.im.2019.103174>
- Selenko, E., Leicht-Deobald, U., Charlier, S.D. and Baldry, C. (2022) "Artificial intelligence and the future of work: A functional-identity perspective", *Group & Organization Management*, Vol 47, No. 4, pp. 737–773. <https://doi.org/10.1177/09637214221091823>

- Simon, H.A. (1960) *The New Science of Management Decision*, New York: Harper & Brothers.
- Venkatesh, V., Thong, J.Y.L. and Xu, X. (2012) "Consumer acceptance and use of information technology: Extending the unified theory of acceptance and use of technology", *MIS Quarterly*, Vol 36, No. 1, pp. 157–178. <https://doi.org/10.2307/41410412>
- Venkatesh, V., Morris, M.G., Davis, G.B. and Davis, F.D. (2003) "User acceptance of information technology: Toward a unified view", *MIS Quarterly*, Vol 27, No. 3, pp. 425–478. <https://doi.org/10.2307/30036540>